

The 2017 KIT IWSLT Speech-to-Text Systems for English and German

Thai-Son Nguyen, Markus Müller, Matthias Sperber, Thomas Zenkel,
Sebastian Stüker and Alex Waibel

Institute for Anthropomatics, Karlsruhe Institute of Technology
Karlsruhe, Germany
`{first.lastname}@kit.edu`

Abstract

This paper describes our German and English *Speech-to-Text* (STT) systems for the 2017 IWSLT evaluation campaign. The campaign focuses on the transcription of unsegmented lecture talks. Our setup includes systems using both the Janus and Kaldi frameworks. We combined the outputs using both ROVER [1] and confusion network combination (CNC) [2] to achieve a good overall performance. The individual subsystems are built by using different speaker-adaptive feature combination (e.g., IMEL with i-vector or bottleneck speaker vector), acoustic models (GMM or DNN) and speaker adaptation (MLLR or fMLLR). Decoding is performed in two stages, where the GMM and DNN systems are adapted on the combination of the first stage outputs using MLLR, and fMLLR.

The combination setup produces a final hypothesis that has a significantly lower WER than any of the individual subsystems. For the English lecture task, our best combination system has a WER of 8.3% on the tst2015 development set while our other combinations gained 25.7% WER for German lecture tasks.

1. Introduction

For many years now, the *International Workshop on Spoken Language Translation* (IWSLT) offers a comprehensive evaluation campaign on spoken language translation. The evaluation is organized in different evaluation tracks covering automatic speech recognition (ASR), machine translation (MT), and the full-fledged combination of the two of them into speech translation systems (SLT). Different from previous years, this year's installment mostly consists of real-life lectures e.g., real university lectures or talks at real symposia.

The goal of the ASR track is the automatic transcription of fully unsegmented lectures. The quality of the resulting transcriptions is measured in word error rate (WER).

This system paper describes our English and German ASR setups with which we participated in the lecture ASR tracks of the 2017 IWSLT evaluation campaign. Similar to previous years' evaluation [3], we used the Janus Recognition Toolkit (JRTk) [4] which features the IBIS single-pass decoder [5] to build several complementary subsystems and combined them with an additional system developed with the

Kaldi toolkit [6]. Our Janus-based systems employ different speaker-adaptive features, acoustic models or speaker adaptation techniques. While the Kaldi-based system applies the same adaptation techniques but employs sequence training and big n-gram language models for rescoring.

The rest of this paper is structured as follows. Section 2 describes the data that our system was trained and tested on. This is followed by Section 3 which provides a description of the acoustic front-ends used in our system and Section 7 which describes our segmentation setup. An overview of the techniques used to build our acoustic models is given in Section 5. We describe the language model used for this evaluation in Section 6. Our decoding strategy and results are then presented in Sections 8 and 9. We conclude the paper with Section 10.

2. Data Resources

2.1. Training Data

Table 1 and Table 2 show the data sources we used for the acoustic model training of our systems. This year we included 80 hours of broadcast news which results in a total of 483 hours for the English systems. For the German systems, we used the same training data as last year.

Source	# Amount
Quaero from 2010 to 2012	200 hours
Broadcast news [7]	80 hours
TED-LIUM v2 [8] excluding disallowed talks	203 hours
Total	483 hours

Table 1: *English acoustic modeling data.*

2.2. Test Data

For this year's evaluation campaign, the evaluation test set "tst2017" as well as the development test sets "tst2015", "tst2013" and "dev2017" were provided for the English and German lecture tasks. All development test sets featured a pre-segmentation provided by the IWSLT organizers. For the

Source	# Amount
Quaero from 2009 to 2012	180 hours
Broadcast news	24 hours
Baden-Württemberg parliament	160 hours
Total	364 hours

Table 2: *German acoustic modeling data.*

evaluation test set, automatic segmentation was required.

3. Feature Extraction

Our systems are built using several different front-ends as previously described in [3] including 40-dimensional log scale mel filterbank (IMEL), 20-dimensional mel frequency cepstral coefficient (MFCC), 20-dimensional minimum variance distortionless response (MVDR) and 14-dimensional tonal (T) features. These features can be augmented with i-vectors (Section 3.2) or bottleneck speaker vectors (Section 3.3) to be directly used for acoustic modeling or fed into deep bottleneck networks (Section 3.1) for extracting bottleneck features. The extracted bottleneck features are then transformed using feature-space maximum likelihood linear regression (fMLLR) and augmented with i-vectors to build speaker-adaptive features (Section 3.4). Our detailed feature extraction pipeline is explained in [9].

3.1. Bottleneck Features

We employed the deep bottleneck architecture described by [10], which consists of a stacked denoising auto-encoder of 4-5 layers each containing 1600-2000 units, followed by a 42 unit bottleneck, a hidden layer and the classification layer. The stacked auto-encoder is first pre-trained layer-wise [11], then the whole network is fine-tuned to discriminate target phoneme states. For the extraction of bottleneck features (BN), the layers after the bottleneck were removed and the output activations of the bottleneck layer were used as BN.

3.2. I-vectors

To extract i-vectors, a full universal background model (UBM) with 2048 mixtures was trained on the training dataset using 20 Mel-frequency cepstral coefficients with delta and delta-delta features appended. The total variability matrices were estimated for extracting 100 dimensional i-vectors. We tuned the size of the i-vectors in a series of preliminary experiments for optimal recognition performance. The UBM model training and i-vector extraction was performed by using the sre08 module from the Kaldi toolkit [6]. I-vectors as well as tonal features were always used in combination with other features.

3.3. Bottleneck Speaker Vectors

In addition to i-vectors, we also used Bottleneck Speaker Vectors (BSVs) [12]. While they serve the same purpose, they are entirely neural network based. We used the same setup as for our hybrid systems, but trained the network to recognize different speakers instead of phonemes using a one-hot encoding of the speaker identities. To extract the BSVs, we used a bottleneck layer as second last layer of the speaker classification network and discarded all layers after this layer after training. For obtaining the final speaker vector, we averaged the output activation of this hidden layer on a per speaker basis.

3.4. Speaker Adaptive Features

To build speaker-adaptive features (SAF) for GMM systems, we first train deep bottleneck network from 11 stacked frames of regular features and i-vectors. The extracted BN features are then spliced for 11 consecutive frames and transformed using Linear Discriminate Analysis (LDA) which are known to make inputs more accurately modeled by GMMs.

The speaker-adaptive features for DNN systems are obtained after transforming BN features using fMLLR transformation and then augmented with i-vectors. The process of fMLLR estimation was performed as traditional approach. During the training, we used the adaptation data of the same speaker and the reference transcriptions to do the alignment, while the same GMMs were used as first-pass systems to generate transcriptions in the testing.

4. Phoneme and Dictionary

For English, we used the CMU dictionary¹. This is the same phoneme set as the one used in last year's systems. It consists of 45 phonemes and allophones. We used 7 noise tags and one silence tag. Missing pronunciations were created using the FESTIVAL [13] Text-to-Speech Engine.

Our German system uses an initial dictionary based on the Verbmobil Phoneset [14]. Missing pronunciations are generated using both MaryTTS [15] and FESTIVAL [13].

5. Acoustic Modeling

5.1. HMM CD-Phone

All GMM and hybrid models classify context-dependent quinphones with three states per phoneme and a left-to-right HMM topology without skip states. The English acoustic models use 8,156 distributions and codebooks derived from decision-tree based clustering of the states of all possible quinphones. The German acoustic models use either 10k or 18k context-dependent states.

¹<http://www.speech.cs.cmu.edu/cgi-bin/cmudict>

5.2. GMM Models

The GMM models are trained by using incremental splitting of Gaussians training (MAS) [16], followed by optimal feature space training (OFS) which is a variant of *semi-tied covariance* (STC) [17] training using a single global transformation matrix. The model is then refined by one iteration of Viterbi training.

For the evaluation, we trained one GMM system using SAF features with MFCC front-ends for the English lecture task.

5.3. Hybrid Models

All the DNN models also share the same architecture which has 5-6 hidden layers with 2000 units per layer. The input of the DNNs are 11 stacked frames of 42-dimensional transformed bottleneck features or 40-dimensional IMEL, with or without combining i-vectors and tonal features. We used the sigmoid activation function for the hidden layers and softmax for the output layer. DNN systems were trained using the cross-entropy loss function to predict context-dependent states. The same training method is applied for all DNNs which includes pre-training with denoising auto-encoders and followed by fine-tuning with back-propagation. We used an exponential schedule to update the learning rate during the neural network training.

This year, we built two DNNs using SA features with different front-ends for the English TED task.

The German setup for the lectures task consists of 4 DNN systems based on different combinations of input features as shown in Table 6. SAF features were used as well.

6. Language Models

6.1. Vocabulary and Kneser-Ney Models

For language model training and vocabulary selection, we used the subtitles of TED talks, or translations thereof, and text data from various sources (see Tables 3 and 4). Text cleaning included tokenization, lowercasing, number normalization, and removal of punctuation. Language model training was performed by building separate language models for all (sub-)corpora using the SRILM toolkit [18] with modified Kneser-Ney smoothing. These were then linearly interpolated, with interpolation weights tuned using held-out data from the TED corpus. For German, we split compounds similarly as in [19].

For the vocabulary selection, we followed an approach proposed by Venkataraman et al. [20]. We built unigram language models using Witten-Bell smoothing from all text sources, and determined unigram probabilities that maximized the likelihood of a held-out TED data set. As our vocabulary, we then used the top 150k words for English, and 300k words for German.

Text corpus	# Words
TED	3.6m
Fisher	10.4m
Switchboard	1.4m
TEDLIUM dataselection	155m
News + News-commentary + -crawl	4,478m
Commoncrawl	185m
GIGA	2323m

Table 3: *English language modeling data.*

Text corpus	# Words
TED	2,685k
News+Newscrawl	1,500M
Euro Language Newspaper	95,783k
Common Crawl	51,156k
Europarl	49,008k
ECI	14,582k
MultiUN	6,964k
German Political Speeches	5,695k
Callhome	159k
HUB5	20k

Table 4: *German language modeling data after cleaning and compound splitting.*

6.2. Feed-forward Neural Language Model

During decoding the probabilities of a feedforward neural network language model were linearly interpolated with the baseline language model. Due to performance considerations, the most recent 40k queries for this language model were cached and we constrained the output vocabulary to the 20k most frequent words which appeared in the text corpora. We used 200 dimensional word embeddings trained with the Skip-gram model [21]. Three words were considered as the context, while the rest of the network consisted of three hidden layers followed by a softmax output layer. The training text consisted of 30M words and was selected based on the text sources listed in Table 4.

7. Automatic Segmentation

In this evaluation, the test set for the ASR track was provided without manual sentence segmentation, thus automatic segmentation of the target data was mandatory. We utilized an approach to automatic segmentation of audio data that is SVM based. This kind of segmentation is using speech and non-speech models, using the framework introduced in [22]. The pre-processing makes use of an LDA transformation on DBNF feature vectors after frame stacking to effectively incorporate temporal information. The SVM classifier is trained with the help of LIBSVM [23]. A 2-phased post-processing is applied for final segment generation.

We generated the segmentations for both English and German using this SVM based segmentation. The parameters for the SVM segmenter were chosen on a per language basis after preliminary experiments.

8. Systems and Combination

Table 5 shows our systems built for the English submission. In the first-pass, we used a GMM and two DNN systems with the acoustic models and 4-gram language model described in Section 5 and Section 6. Their decoded lattices are sent to a consensus decoding system (CNC) to produce combined hypotheses and confidence scores for the adaptation in the second-pass. The GMM system is fully adapted as transitional approach using both feature space adaptation (fMLLR) and model adaptation (MLLR). The DNN systems are adapted by training the DNN acoustic models one more epoch on the adaptation data of each speaker. The adaptation data is obtained by performing alignment of the CNC decoded results with the speaker audio and filtering out the frames with the confidence scores higher than 0.7. All these systems were built using Janus Recognition Toolkit (JRTK) [14].

Beside that we also used the Kaldi toolkit [6] to build a new system with similar feature adaptation techniques. The same train database is used for acoustic modeling and we use the trained 4-gram language model for rescoring.

Our final submission for the English lecture task consists of a ROVER of the Kaldi based system and the adapted systems. The results of the single and adapted systems as well the combined system are presented in Table 5.

9. Results

For the English task, we gained significant improvements over building speaker adaptive features, DNN model adaptation and CNC combination. On the test set “tst2015”, we archived 8.3% WERs.

System	tst2015
GMM(SAF-MFCC)	11.6
DNN(SAF-IMEL)	10.2
DNN(SAF-MFCC)	11.2
CNC	9.4
GMM(SAF-MFCC) adapted	9.3
DNN(SAF-IMEL) adapted	8.8
DNN(SAF-MFCC) adapted	9.3
Kaldi 4-gram LM rescored	9.3
ROVER	8.3

Table 5: Results for English talk task on ‘tst2015’ development set.

In addition to our experiments on these two English

tracks, we also participated in the German lecture task. The results on the “dev2017” test set are shown in Table 6.

System	dev2017
18k DNN(BSV BN-IMEL+T) NNLM	26.7
18k DNN(Mod-M2+IMEL+T)	27.1
10k DNN(SAF-BN-M2+T) NNLM	25.2
10k DNN(SAF-BN-IMEL+T) NNLM	25.7
CNC	25.7

Table 6: Results for German lecture task on ‘dev2017’ development set.

10. Conclusion

In this paper we presented our English and German LVCSR systems, with which we participated in the 2017 IWSLT evaluation. All systems make use of neural network based front-ends, HMM/GMM and HMM/DNN based acoustics models. The decoding set-up of all languages makes extensive use of system combination of single systems obtained by combining different feature extraction front-ends and acoustic models.

11. References

- [1] J. Fiscus, “A post-processing system to yield reduced word error rates: Recognizer Output Voting Error Reduction (ROVER),” in *Proceedings the IEEE Workshop on Automatic Speech Recognition and Understanding*, Santa Barbara, CA, USA, Dec. 1997, pp. 347–354.
- [2] L. Mangu, E. Brill, and A. Stolcke, “Finding consensus in speech recognition: word error minimization and other applications of confusion networks,” *Computer Speech & Language*, vol. 14, no. 4, pp. 373–400, 2000.
- [3] Markus Müller, Thai-Son Nguyen, Matthias Sperber, Kevin Kilgour, Sebastian Stüker, and Alex Waibel, “The 2015 KIT IWSLT Speech-to-Text Systems for English and German,” in *International Workshop on Spoken Language Translation (IWSLT)*, Dec. 2015.
- [4] M. Woszczyna, N. Aoki-Waibel, F. D. Buø, N. Coccaro, K. Horiguchi, T. Kemp, A. Lavie, A. McNair, T. Polzin, I. Rogina, C. Rose, T. Schultz, B. Suhm, M. Tomita, and A. Waibel, “Janus 93: Towards spontaneous speech translation,” in *International Conference on Acoustics, Speech, and Signal Processing 1994*, Adelaide, Australia, 1994.
- [5] H. Soltau, F. Metze, C. Fügen, and A. Waibel, “A one-pass decoder based on polymorphic linguistic context assignment,” in *Automatic Speech Recognition and Understanding, 2001. ASRU’01. IEEE Workshop on*, IEEE, 2001, pp. 214–217.

- [6] D. Povey, A. Ghoshal, G. Boulianne, L. Burget, O. Glembek, N. Goel, M. Hannemann, P. Motlicek, Y. Qian, P. Schwarz, J. Silovsky, G. Stemmer, and K. Vesely, “The kaldi speech recognition toolkit,” in *IEEE 2011 Workshop on Automatic Speech Recognition and Understanding*. IEEE Signal Processing Society, Dec. 2011, iEEE Catalog No.: CFP11SRW-USB.
- [7] D. Graff, “The 1996 broadcast news speech and language-model corpus,” in *Proceedings of the 1997 DARPA Speech Recognition Workshop*, 1996.
- [8] A. Rousseau, P. Deléglise, and Y. Estève, “Enhancing the ted-lium corpus with selected data for language modeling and more ted talks,” in *Proc. of LREC*, 2014, pp. 3935–3939.
- [9] Thai Son Nguyen, Kevin Kilgour, Matthias Sperber, and Alex Waibel, “Improved speaker adaption by combining i-vector and fmlr with deep bottleneck networks,” *SPECOM 2017*. Springer, Cham, 2017.
- [10] J. Gehring, Y. Miao, F. Metze, and A. Waibel, “Extracting deep bottleneck features using stacked autoencoders,” in *Acoustics, Speech and Signal Processing (ICASSP), 2013 IEEE International Conference on*. IEEE, 2013.
- [11] P. Vincent, H. Larochelle, Y. Bengio, and P.-A. Manzagol, “Extracting and composing robust features with denoising autoencoders,” in *The 25th international conference on Machine learning*. ACM, 2008, pp. 1096–1103.
- [12] H. Huang and K. C. Sim, “An investigation of augmenting speaker representations to improve speaker normalisation for dnn-based speech recognition,” in *2015 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2015, pp. 4610–4613.
- [13] A. Black, P. Taylor, R. Caley, and R. Clark, “The festival speech synthesis system,” 1998.
- [14] M. Finke, P. Geutner, H. Hild, T. Kemp, K. Ries, and M. Westphal, “The karlsruhe-verbmobil speech recognition engine,” in *Acoustics, Speech, and Signal Processing, 1997. ICASSP-97., 1997 IEEE International Conference on*, vol. 1. IEEE, 1997, pp. 83–86.
- [15] M. Schröder and J. Trouvain, “The german text-to-speech synthesis system mary: A tool for research, development and teaching,” *International Journal of Speech Technology*, vol. 6, no. 4, pp. 365–377, 2003.
- [16] T. Kaukoranta, P. Fränti, and O. Nevalainen, “Iterative split-and-merge algorithm for VQ codebook generation,” *Optical Engineering*, vol. 37, no. 10, pp. 2726–2732, 1998.
- [17] M. Gales, “Semi-tied covariance matrices for hidden markov models,” *IEEE Transactions on Speech and Audio Processing*, vol. 7, no. 3, pp. 272–281, 1999.
- [18] A. Stolcke, “Srilm-an extensible language modeling toolkit,” in *Seventh International Conference on Spoken Language Processing*, 2002.
- [19] Kevin Kilgour, Christian Mohr, Michael Heck, Quoc Bao Nguyen, Van Huy Nguyen, Evgeniy Shin, Igor Tseyzer, Jonas Gehring, Markus Müller, Matthias Sperber, Sebastian Stüker, and Alex Waibel, “The 2013 KIT IWSLT Speech-to-Text Systems for German and English,” in *International Workshop on Spoken Language Translation (IWSLT)*, Dec. 2013.
- [20] A. Venkataraman and W. Wang, “Techniques for effective vocabulary selection,” in *Proceedings of the 8th European Conference on Speech Communication and Technology*, 2003, pp. 245–248.
- [21] T. Mikolov, K. Chen, G. Corrado, and J. Dean, “Efficient estimation of word representations in vector space,” *arXiv preprint arXiv:1301.3781*, 2013.
- [22] M. Heck, C. Mohr, S. Stüker, M. Müller, K. Kilgour, J. Gehring, Q. Nguyen, V. Nguyen, and A. Waibel, “Segmentation of telephone speech based on speech and non-speech models,” in *Speech and Computer*, ser. Lecture Notes in Computer Science, M. Železný, I. Habernal, and A. Ronzhin, Eds. Springer International Publishing, 2013, vol. 8113, pp. 286–293.
- [23] C.-C. Chang and C.-J. Lin, “LIBSVM: A Library for Support Vector Machines,” *ACM Transactions on Intelligent Systems and Technology*, vol. 2, pp. 27:1–27:27, 2011.