

ROBUST ACOUSTIC AND SEMANTIC CONTEXTUAL BIASING IN NEURAL TRANSDUCERS FOR SPEECH RECOGNITION

Xuandi Fu[†], Kanthashree Mysore Sathyendra[†], Ankur Gandhe, Jing Liu, Grant P. Strimel,
Ross McGowan, Athanasios Mouchtaris

Amazon Alexa

ABSTRACT

Attention-based contextual biasing approaches have shown significant improvements in the recognition of generic and/or personal rare-words in End-to-End Automatic Speech Recognition (E2E ASR) systems like neural transducers. These approaches employ cross-attention to bias the model towards specific contextual entities injected as bias-phrases to the model. Prior approaches typically relied on subword encoders for encoding the bias phrases. However, subword tokenizations are coarse and fail to capture granular pronunciation information which is crucial for biasing based on acoustic similarity. In this work, we propose to use lightweight character representations to encode fine-grained pronunciation features to improve contextual biasing guided by acoustic similarity between the audio and the contextual entities (termed *acoustic biasing*). We further integrate pretrained neural language model (NLM) based encoders to encode the utterance’s semantic context along with contextual entities to perform biasing informed by the utterance’s semantic context (termed *semantic biasing*). Experiments using a Conformer Transducer model on the Librispeech dataset show a 4.62% - 9.26% relative WER improvement on different biasing list sizes over the baseline contextual model when incorporating our proposed acoustic and semantic biasing approach. On a large-scale in-house dataset, we observe 7.91% relative WER improvement compared to our baseline model. On tail utterances, the improvements are even more pronounced with 36.80% and 23.40% relative WER improvements on Librispeech rare words and an in-house testset respectively.

Index Terms— contextual biasing, attention, RNN-T, Conformer, end-to-end ASR, neural transducers, personalization

1. INTRODUCTION

E2E ASR systems are gaining popularity due to their monolithic nature and ease of training, making them promising candidates for deployment in commercial voice assistants (VAs). However, E2E models typically rely on word-piece vocabularies causing rare entities to often decompose into target sequences that are infrequent in training data making it difficult for the model to recognize them correctly [1–5]. To provide the best user experience, it is important for a voice assistant to incorporate each user’s custom environment and preferences, and use them to improve recognition of personalized requests like “call [Contact Name]”, and “play my [Play List] on spotify”. It is also desirable to recognize other types of rare-phrases like trending words (eg. “coronavirus”) or rare-words that appear on the screen (eg. rare-phrases from a wikipedia page displayed on the VA’s screen or rare-phrases from a previous TTS response). Contextual biasing is widely used to inject such lexical contexts to bias

the model’s predictions towards them. Examples of lexical contexts include user defined terms (eg. user playlists, user’s contacts) and generic rare-phrases (on-screen rare-words, trending words etc.).

Neural contextual biasing methods for E2E ASR models broadly fall into two categories – graph fusion approaches (eg. Deep Shallow fusion, Trie and WSFT-based approaches) [6–8] and fully-neural attention-based approaches [2, 3, 9–12]. In the first category, [1, 13] introduced deep shallow fusion and trie-based contextual biasing with deep personalized LM fusion, respectively. More recently, fully-neural attention-based contextual biasing methods are becoming popular due to their improved rare-word recognition and ease of integration with E2E neural inference engines. Neural contextual biasing for LAS models was explored in [2, 14] where bias phrases and/or contextual entities are encoded via BiLSTM encoders, and the model is made to bias toward them via a location-aware attention mechanism [15–17]. For RNN-T models, [3, 9, 10] introduced fully-neural attention-based contextual biasing. [9, 10] introduce contextual biasing for both the audio encoder and the text-based prediction network in transformer transducers (T-T) and conformer transducers (C-T). [10] explored contextual adaptation of pretrained RNN-T and C-T models using contextual adapters that are faster, cheaper and data-efficient.

To embed the contextual entities, prior approaches relied on a text encoder (typically LSTM-based) that used grapheme-based subword/sentence-piece tokenizations [2, 3, 10, 12]. However, two acoustically similar words may have completely different and non-overlapping subword tokenizations. For instance, the words “seat” and “meat”, although acoustically similar, have very different subword tokenizations [sea, t] and [meat]. Assuming “seat” is a new incoming OOV word, it may be incorrectly predicted as “sit” (which could be seen more often in training data). A character-level encoder may capture the similarities in pronunciations in “seat” and “meat” better. Joint grapheme and phoneme context embedding has also been proposed by [16, 17] to incorporate pronunciation features. However, this requires a pronunciation lexicon or a grapheme-to-phoneme network. Further, a single word can have multiple pronunciations, and can result in longer lists of biasing-phrases. To overcome these issues, we use character-embedding models that not only capture fine-grained pronunciation information, but are also lightweight owing to their small character vocabularies, and require fewer model parameters.

Approaches have also been proposed to leverage semantic context features with language models to improve rare word recognition in ASR systems [18, 19]. In [3, 9, 10], the transducer prediction network output is directly used as query in the prediction-network-side contextual biasing layer. However, the ability of the prediction network in encoding semantic context is limited by its training data, where the semantic context for long-tail utterances and rare words is not well trained [20]. This necessitates the integration of a pretrained language model (PLM) trained on ample text-only data that is capable

[†]Equal Contribution

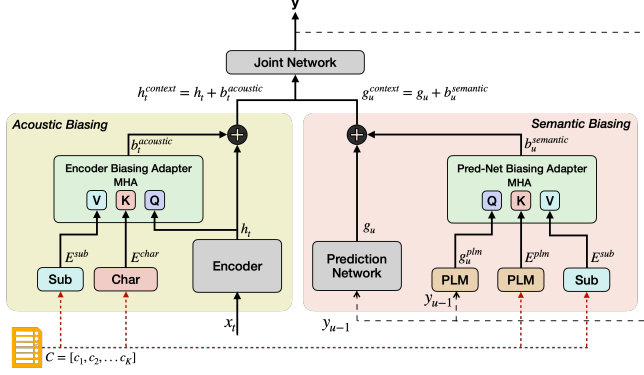


Fig. 1: The proposed acoustic and semantic contextual biasing model. *Char*, *Sub* and *PLM* denote the character-, subword- and PLM-encoder respectively. Same color denotes shared parameters.

of encoding semantics better. Using representations from the PLM as the query to the prediction-network-side biasing layer can help perform better semantic biasing. More generally, semantic biasing addresses the problem of “given the previously decoded word-pieces, what contextual entity from the list is most likely to occur next?”. It ensures semantic coherence between the word-pieces predicted by the model and the selected contextual entity.

In this paper, we propose to train a character-based context representation on the encoder-side to capture generalizable acoustic context features useful for biasing based on acoustic similarity between the contextual entity and the incoming audio. We further introduce neural language model encoders to encode the utterance context and lexical contexts used for biasing based on semantics and meaning of the utterance, thus generalizing to open domain use cases.

2. METHODOLOGY

Neural transducer models typically consist of an encoder network (H), a prediction-network (P) and joint network (J) [21]. The inputs to the encoder networks are audio features of length T , $\mathbf{X} = (x_0, x_1, \dots, x_T)$. The encoder network H (typically stacked recurrent neural networks or conformer blocks), encodes the audio frame x_t into D -dimensional intermediate representation $h_t = H(x_t)$. The prediction-network encodes the sequence of previously predicted non-blank word-pieces $\mathbf{Y} = (y_0, y_1, \dots, y_{u-1})$ into hidden states g_u . The joint network fuses the encoder hidden states at time step t and prediction-network hidden states at step u to compute the probability of next word-piece as, $z_{t,u} = J(h_t + g_u)$. J is a series of feed-forward layers with optional non-linear activations. The transducer model is optimized with the RNN-T loss that computes the alignment probability using the forward-backward algorithm [21]. In order to contextualize the transducer model and help recognition of contextual entities, a list of K bias-phrases, $\mathbf{C} = \{c_1, c_2, \dots, c_K\}$, are encoded and injected to the transducer model [9, 10]. The encoder-side biasing and prediction network-side biasing adapters are responsible for biasing towards the correct contextual entity by learning acoustic and semantic context-aware encoder and prediction network representations, $h_t^{context}$ and $g_u^{context}$ respectively, as shown in Figure 1.

2.1. Acoustic Biasing

The encoder biasing module relies on acoustic similarity between current audio features and the biasing list to attend towards the most acoustically similar contextual entity. The attention-based biasing

adapter computes the cross-modal (from audio and text representations) similarity by using encoder state as the query, and the biasing list as keys and values. We use character-encoders to encode the bias-phrases and use it as the key for attention computation in the encoder biasing adapter, linking the acoustic hidden states h_t and the external context. The proposed encoder contextual adapter consists of three components, a character-level key encoder, a subword-level value encoder and a cross-attention based biasing adapter for incorporating contextual information into the encoder hidden states.

2.1.1. Character Key Encoder

The character key encoder encodes each contextual entity, c_k , into character context embeddings e_k^{char} . Each contextual entity, c_k , is first split into a sequence of characters c_k^{char} and passed through an embedding layer followed by stacked Bi-LSTM layers, the final state of BiLSTM outputs is taken as the key embeddings, denoted as $E^{char} = [e_1^{char}, e_2^{char}, \dots, e_K^{char}] \in \mathbb{R}^{K \times D^c}$ where $e_k^{char} = BiLSTM(Embedding(c_k^{char}))$ and D^c is the embedding size for the character embeddings.

2.1.2. Subword Value Encoder

The subword encoder is used for generating the value for the attention-based biasing layer, which informs the model’s subword predictions (hence uses sub-word embeddings). We use the same sentence-piece model as the C-T model to be consistent with the output unit of the transducer model [22, 23]. The tokenized subword sequences are passed through an embedding layer followed by stacked BiLSTM layers, and the final state of the BiLSTM is forwarded as the subword-level context embedding denoted as $E^{sub} = [e_1^{sub}, e_2^{sub}, \dots, e_K^{sub}]$ and each embedding is a D^s dimensional vector.

2.1.3. Encoder Biasing Adapter (EBA) for Acoustic Biasing

The EBA learns to bias toward relevant context entities by relying on the acoustic similarity between the audio and the biasing list [10]. The multi-head attention (MHA) module takes in audio/encoder hidden states h_t as the *query* and a list of character context embedding for bias-phrases E^{char} as *keys* to compute per-entity attention scores a_t^k (Equation 1) for $t \in [0, T]$. The *query* and *keys* are first projected via $\mathbf{W}^k$ and $\mathbf{W}^q$ correspondingly. The attention score a_t^k , computed by the scaled dot product attention mechanism, measures the similarity between the audio features at time step t and bias phrase c_k . The biasing vector $b_t^{acoustic}$, is computed by the attention-weighted sum of the projected *value* embeddings (projected via $\mathbf{W}^v$) and added back to h_t to obtain context-aware features $h_t^{context}$ (Equation 2).

$$a_t^k = \text{Softmax}\left(\frac{\mathbf{W}^k e_k^{char} \mathbf{W}^q h_t}{\sqrt{d}}\right) \quad (1)$$

$$b_t^{acoustic} = \sum_k a_t^k \mathbf{W}^v e_k^{sub}; \quad h_t^{context} = h_t + b_t^{acoustic} \quad (2)$$

2.2. Semantic Biasing

One limitation of E2E models is that the decoder/prediction-network has to be trained on transcribed audio data which covers limited carrier phrases or semantic utterance contexts as compared to a large-scale text-only corpus. Since acoustic biasing learns similarity between context embedding and audio frames, it tends to give equal attention to homophones and cannot differentiate between homophones. In these cases, a strong indication from the utterance’s semantic context can help the model select the right entity (among words with similar pronunciations) to bias towards. For this, we introduce a

Acoustic Biasing Only	Acoustic Biasing + Semantic Biasing
show me that janie cow	show me that giant cow
what's the simple recipe for recipe again	what's the simple syrup recipe again
what is laser beam latest video	what is laser beam's latest video
find richard samsung tv	find richard's samsung tv
computer play station extreme my radio from tunein	computer play station extreme on my radio from tunein
what is cognag made out of	what is cognac made out of
order gillette sport triumph triumph deodorant for men	order gillette sport triumph deodorant for men
play songs by francisco cabrel	play songs by francis cabrel
play nobody's better by petit wap	play nobody's better by fetty wap
read the book owen by kevin hunks	read the book owen by kevin henkes

Fig. 2: Examples of semantic improvements (green), syntactic improvements (red), spelling corrections (blue) and external knowledge integration (yellow) from PLM Models.

pretrained language model (PLM)-based query and key encoder for prediction network biasing.

2.2.1. PLM-based Query-Key Encoder

The PLM-based encoder (henceforth referred to as PLM-encoder) encodes the previously predicted word pieces $(y_0, y_1, \dots, y_{u-1})$ to produce an utterance semantic context vector $g_u^{plm} \in \mathbb{R}^{1 \times D^p}$, where D^p is PLM embedding size. The same PLM-encoder is also used to encode the bias-phrases, resulting in a shared embedding space between the *queries* and the *keys*. The bias phrases in $\mathcal{C}$ are first tokenized into subword units using the same sentencepiece model as the transducer, and encoded with the PLM to obtain $E^{plm} \in \mathbb{R}^{K \times D^p}$.

2.2.2. Pred-Net Biasing Adapter (PNBA) for Semantic Biasing

The PNBA adopts a similar scaled-dot-product attention as the encoder side (Equation 1) to compute the *semantic similarity* between the utterance context and the bias phrases. The *query* is the PLM-generated semantic embedding g_u^{plm} encoding the utterance semantic context, *key* is the PLM-generated context embedding E^{plm} . For *value*, we use the subword value embedding E^{sub} , shared with the encoder-side biasing adapter. The semantic prediction network biasing vector, $b_u^{semantic}$ is then computed similar to Equation 2 for u , and fused with the prediction network hidden states as $g_u^{context} = g_u + b_u^{semantic}$. The generated context-aware hidden states, $h_t^{context}$ and $g_u^{context}$, are passed to the joint network to predict the probability for the next word piece, $z_{t,u}^{context} = J(h_t^{context} + g_u^{context})$. The biasing adapters are trained with the RNN-T loss, through the contextual-adaptation approach proposed by [10].

3. EXPERIMENTAL RESULTS

3.1. Datasets and Evaluation Metrics

Our experiments are conducted on LibriSpeech [25] and an in-house large-scale voice assistant dataset. For the LibriSpeech datasets, a global rare word list is constructed by removing the top 5000 common words from the training vocabulary as in [19]. During training, $K = 50$ randomly sampled distractors from the global rare words list are added to the biasing list, along with the correct words and the special $\langle \text{no-bias} \rangle$ token [2, 10]. The biasing list is formulated in a similar fashion during evaluation and includes the correct entity and $\langle \text{no-bias} \rangle$. We further evaluate with larger biasing lists for $K = \{50, 100, 500, 1000\}$.

The in-house dataset consists of 18k hours of de-identified audio data containing rare-words from multiple domains, such as *song names*, *location names*, *device names*, *playlist names* and other tail entities. The models are evaluated on four testsets with $K = 100$ biasing list size. *General*: test utterances sampled follow the original training

data distribution (contain both rare and non-rare utterances). *Rare Word*: test utterances with at least one rare word (i.e. not in top-1000 most frequent words in training data). *Zero-shot Rare Words*: testset containing tail utterances with rare words that are not in training data. *Knowledge*: To test the semantic biasing performance of the models, we also prepare a test set with diverse semantic contexts, by pooling utterances from domains like *Wikipedia*, *Knowledge* and *Books*. We also test the performance of our models when supplying personalized contexts on datasets – *Proper Names* and *Devices*.

We evaluate the model performance on four metrics: (1) *WER*: overall word error rate on the whole test set; (2) *R-WER*: word error rate for rare words (RW) requiring biasing (3) *NR-WER*: word error rate for non-rare words that do not require biasing (4) *ZSR-WER*: word error rate for zero-shot rare words i.e. words not seen during training. For the in-house dataset, results are presented as relative error rate reductions¹ (*WERR*, *R-WERR*, *NR-WERR* and *ZSR-WERR*)².

3.2. Baseline and Model Configuration

The proposed methods are compared against the contextual biasing adapter baseline using only the subword encoder to encode contextual entities [10]. We follow the same adapter-type training strategy as in [10] to first train the core transducer models and then train the adapters on *Mixed* dataset with the base transducer parameters frozen to avoid degradation on recognizing general words.

ASR model. The input audio features are 64-dimensional LFBF features extracted every 10 ms with a window size of 25 ms. The features of each frame are then stacked with the left two frames, followed by a downsampling factor 3 to achieve a low frame rate, resulting in 192-dimensional features. For training the biasing layer, we used *Mixed* dataset of utterances that contain rare words and general utterances containing common words to teach the model to learn the representations for context entities and the $\langle \text{no-bias} \rangle$ token. The contextual biasing adapters are trained with streaming conformer-transducer (C-T) architecture to collaboratively decode the word-pieces. For LibriSpeech, we use 14 conformer blocks [24] for the audio encoder. The hidden state size d_{model} , the projection dimension d_{ff} and the number of attention heads h for the conformer encoder are 256, 1024 and 4 correspondingly. For the prediction-network, the transcriptions are first tokenized by 2500 word-piece SentencePiece tokenizer [23], and then passed through a 256-dim embedding layer followed by a 1-layer LSTM with 640 units and output size of 512-dim. For the in-house dataset, the audio encoder has 12 conformer blocks with $d_{model} = 512$, $d_{ff} = 512$ and $h = 4$, and output size is 512-dim. The prediction network encodes word-pieces with a 512-dim embedding layer and a 2-layer LSTM with 736 units and the output is projected to 512-dim.

Char-Subword Encoder Biasing Model For the char-subword biasing model for Librispeech, the context entities are tokenized with a 38-character tokenizer then embedded with a 64-dim embedding layer followed by one 64-unit BiLSTM layer. For subword representations, the contextual entities are tokenized using a 2500 sentence-piece tokenizer and encoded with a 64-dim embedding layer followed by 64-unit BiLSTM layer. The query, key and values for the multi-head attention layers in the acoustic biasing layer are projected to 128-dim tensors. For the in-house dataset, same-sized character-level context encoder and subword context encoder are used whereas the contextual entities are tokenized by a 4000 sentence-piece tokenizer

¹Due to internal company policies, we are unable to report absolute WER metrics on in-house data

²Given a model A's WER (WER_A) and a baseline B's WER (WER_B), the WERR of A over B is computed as $WERR = (WER_B - WER_A) / WER_B$.

Table 1: Librispeech. Absolute WER for the models for $K = \{50, 100, 500, 1000\}$. Inputs to the Encoder biasing adapters in the form (Key, Value), and Pred-Net biasing adapters in the form (Query, Key, Value). ‘Sub’, ‘Char’, ‘PLM’ and ‘PredNet’ are abbreviations to denote sub-word embeddings E^{sub} , character embeddings E^{char} , PLM embeddings E^{plm} and the Prediction network states g_u respectively.

Model Type	Encoder Biasing Adapter Inputs (K, V)	Pred-Net Biasing Adapter Inputs (Q, K, V)	test-clean				test-other			
			K=50	K=100	K=500	K=1000	K=50	K=100	K=500	K=1000
C-T [24]	-	-	6.08				14.01			
Baseline (B1) [10]	(Sub, Sub)	-	4.70	4.75	5.05	5.19	12.08	12.21	12.75	13.05
Char-I	(Char, Char)	-	4.98	5.06	5.30	5.43	12.67	12.70	13.20	13.55
Char-II	(Char, Sub)	-	4.56	4.63	4.82	5.08	11.80	12.06	12.51	12.89
Char-Subword	(Char, Sub)	(PredNet, Sub, Sub)	4.27	4.31	4.67	4.95	11.16	11.34	12.21	12.64
Subword-PLM	(Sub, Sub)	(PredNet, Sub, Sub)	4.22	4.32	4.63	4.95	11.01	11.22	12.11	12.72
Char-PLM	(Char, Sub)	(PLM, PLM, Sub)	4.10	4.11	4.52	4.83	11.01	11.21	12.09	12.60

Table 2: In-house datasets. WERR on a variety of in-house testsets. The inference contextual entities (C) is one of *rare-words (RW)*, *proper-names (PN)* or *device-names (D)*. $K=100$.

Model	General	RW	ZSRW	Knowledge	Proper Names	Devices
	C = (RW)	C = (RW)	C = (RW)	C = (RW)	C = (PN)	C = (D)
Baseline (B1)	0.00%	0.00%	0.00%	0.00%	0.00%	0.00%
Char-I	-4.11%	-9.54%	-17.45%	-5.18%	-13.78%	+3.10%
Char-II	+0.91%	+1.19%	+5.52%	+2.38%	+3.24%	+1.65%
Char-Subword	+5.18%	+8.48%	+16.40%	+6.63%	+12.24%	+5.58%
Subword-PLM	+7.00%	+12.58%	+15.80%	+9.02%	+8.91%	+3.10%
Char-PLM	+7.91%	+14.97%	+23.40%	+11.40%	+13.13%	+4.75%

Table 3: Few-shot generalization on Librispeech (test-clean). WER for n -shot rare words; $n \in \{0, 1, 5, 10, 20, 100\}$. $K=100$.

Model	ZSR-WER						R-WER	NR-WER
	0-shot	≤1-shot	≤5-shot	≤10-shot	≤20-shot	≤100-shot		
C-T	113.6	103.9	85.4	70.0	54.5	27.7	24.6	3.8
Baseline (B1)	91.2	81.6	60.1	46.2	33.5	15.2	13.1	3.9
Char-II	85.5	75.2	55.2	42.1	30.4	13.7	12.0	3.8
Char-Subword	79.2	67.7	48.8	37.0	26.0	11.5	10.0	3.7
Subword-PLM	81.9	67.8	47.9	36.0	25.5	11.0	9.6	3.7
Char-PLM	71.0	58.8	42.3	31.5	22.0	9.6	8.3	3.6
WERR	+22.1%	+27.9%	+29.6%	+31.8%	+34.3%	+36.8%	+36.6%	+7.7%
Word Count	331	461	814	1154	1772	4844	5752	46815

for obtaining the subword representations. For training, we use the Adam optimizer with learning rate $5e-4$ and $1e-4$ for the LibriSpeech and the in-house model respectively.

Pretrained Language Model (PLM) encoder: For LibriSpeech, the PLM-encoder has a 256-dim embedding layer followed by a 2-layer LSTM (256 units each), trained on the text-only Wikitext-103 corpus [26]. For in-house data, a domain-general NLM (embedding-dim 512, lstm-size 512) was trained on 80 million utterances with live audio from anonymized user interactions. To improve coverage on rare entities, we also added 25 million entities from artist and song name catalogs, and 8 million entities from place name catalogs.

3.3. Results

Table 1,2 show WER for the Librispeech and in-house datasets. On Librispeech, the proposed acoustic-semantic biasing model outperforms the subword-embedding baseline [10] for varying biasing list sizes(+9.26% and +4.62% WERR for test-clean $K=100$ and 1000). For the in-house dataset, it shows improvements on both *General* (+7.91% WERR) and *Rare-Word* testsets (+14.97% WERR).

Improved generalization abilities. Character key encoder along with PLM-based prediction network biasing shows significant improvements over baseline (B1) on recognizing zero-shot and few-shot rare words. On Librispeech, WER on 0-shot rare words, improved from 91.2% to 85.5% for *Char-II* when compared with B1, and fur-

Table 4: Accuracy of the attention mechanism in selecting the correct entity for biasing, for different Query-Key (Q -, K -) pairs for PNBA.

Models	AA Relative(%)	AAS Relative(%)
(Q-PredNet, K-Sub)	0.00	0.00
(Q-PredNet, K-PLM)	+2.25	+13.80
(Q-PLM, K-PLM)	+7.49	+16.55

ther improved to 71.0% with PLM-based prediction network biasing (Table 3). The relative improvement becomes more pronounced as training samples for the rare words increase (+36.8% WERR for ≤ 100 -shot). On the in-house ZSRW dataset, the *Char-II* resulted in +5.52% WERR. To tease out the improvements from PLM integration, we compared *Char-PLM* (PLM-based query-key) with *Char-Subword* model (g_u as query, *Sub* for key). On ZSRW testset, WERR goes from 16.40% (for *Char-Subword*) to 23.40% for *Char-PLM*(Table 2).

Importance of PLM-based Prediction Network biasing. The PLM-based query-key biasing layers demonstrate improved recognition in terms of semantic, syntactic, spelling improvements and external knowledge integration. Some qualitative examples are shown in Figure 2. With the external knowledge learned from larger text corpus, the PLM-based biasing layers are able to correctly attend over the artist name “*francis cabrel*”. On the *Knowledge* in-house testset containing tail/rare words from multiple domains, the PLM-based model has shows 11.40% WERR compared over B1.

Evaluation of attention mechanism. To compare the ability of the proposed semantic biasing model to attend over the correct bias phrase (vs. over distractors) based on semantic similarity, we compute two additional metrics – 1) *Attention Accuracy (AA)(%)*: % of utterances where the model attends over the correct contextual entity, i.e. has the highest attention score for the correct entity. 2) *Avg. Attention Score (AAS)(%)*: the average attention score for the correct bias phrase assigned by the attention models. Comparisons are made on the in-house ZSRW testset (Table 4 - reported as relative improvement). The PLM-based biasing model((*Q-PLM, K-PLM*)) shows 7.49% relative and 16.55% relative improvement in AA and AAS respectively, when compared with using prediction-network output as query. Thus, demonstrating its ability to leverage the utterance semantic context in selecting the entities to attend over.

4. CONCLUSION

In this work, we propose to improve conformer-transducer biasing models with character-based acoustic biasing and PLM-based semantic biasing model to link the utterance context with the correct biasing entities. Through experiments on LibriSpeech and large-scale in-house datasets, we demonstrated that the proposed acoustic-semantic biasing model can encode better context features and learns a superior biasing model that is guided by acoustic and semantic similarity. It also shows significant improvements when generalizing to tail rare words and open-domain use cases compared to the baseline.

5. REFERENCES

- [1] Duc Le, Gil Keren, Julian Chan, Jay Mahadeokar, Christian Fuegen, and Michael L Seltzer, “Deep shallow fusion for rnn-t personalization,” in *2021 IEEE Spoken Language Technology Workshop (SLT)*. IEEE, 2021, pp. 251–257.
- [2] Golan Pundak, Tara N Sainath, Rohit Prabhavalkar, Anjuli Kannan, and Ding Zhao, “Deep context: end-to-end contextual speech recognition,” in *2018 IEEE spoken language technology workshop (SLT)*. IEEE, 2018, pp. 418–425.
- [3] Mahaveer Jain, Gil Keren, Jay Mahadeokar, Geoffrey Zweig, Florian Metze, and Yatharth Saraf, “Contextual rnn-t for open domain asr,” *Proc. Interspeech 2020*, pp. 11–15, 2020.
- [4] Tara N Sainath, Rohit Prabhavalkar, Shankar Kumar, Seungji Lee, Anjuli Kannan, David Rybach, Vlad Schogol, Patrick Nguyen, Bo Li, Yonghui Wu, et al., “No need for a lexicon? evaluating the value of the pronunciation lexica in end-to-end models,” in *2018 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2018, pp. 5859–5863.
- [5] Tony Bruguier, Fuchun Peng, and Françoise Beaufays, “Learning personalized pronunciations for contact names recognition,” 2016.
- [6] Ding Zhao, Tara N Sainath, David Rybach, Pat Rondon, Deepti Bhatia, Bo Li, and Ruoming Pang, “Shallow-fusion end-to-end contextual biasing,” 2019.
- [7] Yanzhang He, Tara N Sainath, Rohit Prabhavalkar, Ian McGraw, Razi Alvaraz, Ding Zhao, David Rybach, Anjuli Kannan, Yonghui Wu, Ruoming Pang, et al., “Streaming end-to-end speech recognition for mobile devices,” in *2019 IEEE ICASSP*. IEEE, 2019, pp. 6381–6385.
- [8] Aditya Gourav, Linda Liu, Ankur Gandhe, Yile Gu, Guitang Lan, Xiangyang Huang, Shashank Kalmane, Gautam Tiwari, Denis Filimonov, Ariya Rastrow, et al., “Personalization strategies for end-to-end speech recognition systems,” in *2021 IEEE ICASSP*. IEEE, 2021, pp. 7348–7352.
- [9] Feng-Ju Chang, Jing Liu, Martin Radfar, Athanasios Mouchtaris, Maurizio Omologo, Ariya Rastrow, and Siegfried Kunzmann, “Context-aware transformer transducer for speech recognition,” in *2021 IEEE Automatic Speech Recognition and Understanding Workshop (ASRU)*. IEEE, 2021, pp. 503–510.
- [10] Kanthashree Mysore Sathyendra, Thejaswi Muniyappa, Feng-Ju Chang, Jing Liu, Jinru Su, Grant P Strimel, Athanasios Mouchtaris, and Siegfried Kunzmann, “Contextual adapters for personalized speech recognition in neural transducers,” in *ICASSP 2022-2022 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2022, pp. 8537–8541.
- [11] Minglun Han, Linhao Dong, Zhenlin Liang, Meng Cai, Shiyu Zhou, Zejun Ma, and Bo Xu, “Improving end-to-end contextual speech recognition with fine-grained contextual knowledge selection,” in *ICASSP 2022-2022 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2022, pp. 8532–8536.
- [12] Tsendsuren Munkhdalai, Khe Chai Sim, Angad Chandorkar, Fan Gao, Mason Chua, Trevor Strohman, and Françoise Beaufays, “Fast contextual adaptation with neural associative memory for on-device personalized speech recognition,” in *2022 IEEE ICASSP*, 2022, pp. 6632–6636.
- [13] Ding Zhao, Tara N. Sainath, David Rybach, Pat Rondon, Deepti Bhatia, Bo Li, and Ruoming Pang, “Shallow-Fusion End-to-End Contextual Biasing,” in *Proc. Interspeech 2019*, 2019, pp. 1418–1422.
- [14] William Chan, Navdeep Jaitly, Quoc V. Le, and Oriol Vinyals, “Listen, attend and spell: A neural network for large vocabulary conversational speech recognition,” in *ICASSP*, 2016.
- [15] Jan Chorowski, Dzmitry Bahdanau, Dmitriy Serdyuk, Kyunghyun Cho, and Yoshua Bengio, “Attention-based models for speech recognition,” *arXiv preprint arXiv:1506.07503*, 2015.
- [16] Zhehuai Chen, Mahaveer Jain, Yongqiang Wang, Michael L Seltzer, and Christian Fuegen, “Joint grapheme and phoneme embeddings for contextual end-to-end asr,” 2019.
- [17] Antoine Bruguier, Rohit Prabhavalkar, Golan Pundak, and Tara N Sainath, “Phoebe: Pronunciation-aware contextualization for end-to-end speech recognition,” in *ICASSP 2019-2019 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2019, pp. 6171–6175.
- [18] Ashish Shenoy, Sravan Bodapati, Monica Sunkara, Srikanth Ronanki, and Katrin Kirchhoff, “Adapting long context nlm for asr rescoring in conversational agents,” *arXiv preprint arXiv:2104.11070*, 2021.
- [19] Duc Le, Mahaveer Jain, Gil Keren, Suyoun Kim, Yangyang Shi, Jay Mahadeokar, Julian Chan, Yuan Shangguan, Christian Fuegen, Ozlem Kalinli, et al., “Contextualized streaming end-to-end speech recognition with trie-based deep biasing and shallow fusion,” *arXiv preprint arXiv:2104.02194*, 2021.
- [20] Mohammadreza Ghodsi, Xiaofeng Liu, James Apfel, Rodrigo Cabrera, and Eugene Weinstein, “Rnn-transducer with stateless prediction network,” in *ICASSP 2020*. IEEE, 2020, pp. 7049–7053.
- [21] Alex Graves, “Sequence transduction with recurrent neural networks,” *arXiv e-prints*, pp. arXiv–1211, 2012.
- [22] Taku Kudo, “Subword regularization: Improving neural network translation models with multiple subword candidates,” in *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 2018, pp. 66–75.
- [23] Taku Kudo and John Richardson, “Sentencepiece: A simple and language independent subword tokenizer and detokenizer for neural text processing,” in *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, 2018, pp. 66–71.
- [24] Anmol Gulati, James Qin, Chung-Cheng Chiu, Niki Parmar, Yu Zhang, Jiahui Yu, Wei Han, Shibo Wang, Zhengdong Zhang, Yonghui Wu, et al., “Conformer: Convolution-augmented transformer for speech recognition,” *Proc. Interspeech 2020*, pp. 5036–5040, 2020.
- [25] Vassil Panayotov, Guoguo Chen, Daniel Povey, and Sanjeev Khudanpur, “Librispeech: an asr corpus based on public domain audio books,” in *2015 IEEE international conference on acoustics, speech and signal processing (ICASSP)*. IEEE, 2015, pp. 5206–5210.
- [26] Stephen Merity, Caiming Xiong, James Bradbury, and Richard Socher, “Pointer sentinel mixture models,” in *International Conference on Learning Representations*, 2017.