

A Randomized Link Transformer for Diverse Open-Domain Dialogue Generation

Jing Yang Lee¹, Kong Aik Lee², Woon Seng Gan³

School of Electrical and Electronic Engineering, Nanyang Technological University^{1,3}

Institute for Infocomm Research, A*STAR²

jingyang001@e.ntu.edu.sg¹, lee_kong_aik@i2r.a-star.edu.sg², ewsgan@ntu.edu.sg³

Abstract

A major issue in open-domain dialogue generation is the agent’s tendency to generate repetitive and generic responses. The lack in response diversity has been addressed in recent years via the use of latent variable models, such as the Conditional Variational Auto-Encoder (CVAE), which typically involve learning a latent Gaussian distribution over potential response intents. However, due to latent variable collapse, training latent variable dialogue models are notoriously complex, requiring substantial modification to the standard training process and loss function. Other approaches proposed to improve response diversity also largely entail a significant increase in training complexity. Hence, this paper proposes a Randomized Link (RL) Transformer as an alternative to the latent variable models. The RL Transformer does not require any additional enhancements to the training process or loss function. Empirical results show that, when it comes to response diversity, the RL Transformer achieved comparable performance compared to latent variable models.

1 Introduction

Open-domain dialogue generation refers to the task of generating coherent, natural and human-like dialogue given solely the dialogue context (also known as the dialogue history). The development of open-domain dialogue agents that can engage humans in seamless general conversation (or chit-chat) is one of the main objectives of conversational AI. Currently, however, agents display a tendency to generate repetitive and generic dialogue responses, which negatively impact both naturalness and contextual coherence.

Recently, researchers have turned to latent variable models, specifically the Conditional Variational Auto Encoder (CVAE) (Sohn et al., 2015), to address this issue (Yang et al., 2021; Gao et al., 2019; Zhao et al., 2017; Cao and Clark, 2017).

In addition to open-domain dialogue generation, latent variable models have also been applied to related tasks such as personalized dialogue (Lee et al., 2022; Wu et al., 2020; Song et al., 2019), empathetic dialogue (Li et al., 2021, 2020b,a; Zhou and Wang, 2018), and topical dialogue generation (Wang et al., 2020). These works have generally involved modelling the potential dialogue response intents as a latent Gaussian prior, which is typically generated by a Multi-Layer Perceptron (MLP). During inference, a latent instance is sampled from the generated Gaussian prior via the reparameterization trick and fed to the decoder. Stochasticity is induced during the response generation via such random sampling process. However, even though latent variable models are effective at improving response diversity, they are notoriously hard to train primarily due to the Kullback-Liebler (KL) vanishing problem. Usually, this problem is addressed via KL annealing or incorporating a Bag-of-Words loss. While KL annealing requires tuning the weighting hyperparameter β , attaining the Bag-of-Words loss involves defining an additional task of predicting the response bag-of-words. This results in additional complexity during training.

Several other approaches to promoting response diversity which involve introducing an alternate loss function such as the Maximum Mutual Information (MMI) objective (Li et al., 2016a), the Inverse Token Frequency (ITF) objective (Nakamura et al., 2018), and the Inverse N-gram Frequency (INF) objective (Ueyama and Kano, 2020), require considerable additional computation steps. On the other hand, adversarial learning-based (Li et al., 2017a) and embedding augmentation-based (Cao et al., 2021) approaches require extensive modifications to the standard training process, resulting in a significant increase in training complexity.

Hence, inspired by randomization-based neural networks (Suganthan and Katuwal, 2021) and the Random Vector Functional Link (RVFL) neural

network (Pao and Takefuji, 1992) in particular, we introduce a novel Randomized Link (RL) Transformer. An alternative to latent variable models which does not require any modifications or enhancements to the standard training process or loss function. In other words, the RL Transformer can be trained via standard gradient descent solely on the standard negative log-likelihood loss. Experimental results on the DailyDialog (Li et al., 2017b) and EmpatheticDialogues (Rashkin et al., 2019) corpora show that our RL Transformer successfully improves the diversity of the generated response, achieving comparable response diversification relative to latent variable approaches. In addition, compared to the latent variable models, responses generated by our RL Transformer are noticeably more fluent and contextually coherent.

The remainder of this paper is organized as follows: Section 2 provides additional background information regarding open-domain dialogue generation and RVFL neural networks; Section 3 describes the proposed RL Transformer in detail; Section 4 provides details regarding our implementations and experiments; Section 5 presents the experimental results as well as our analysis of the results; Section 6 concludes the paper.

2 Related Work

2.1 Neural Open-domain Dialogue Generation

In recent years, generative neural models have been commonly applied to the task of open-domain dialogue generation. Influenced by advances in machine translation, popular approaches to this task featured a sequence-to-sequence (seq2seq) architecture (Sutskever et al., 2014). Recurrent neural networks such as the Long Short Term Memory (LSTM) and the Gated Recurrent Unit (GRU) have been often utilized in both the encoder and decoder of a seq2seq model (Shang et al., 2015; Sordoni et al., 2015). More recently, Transformer-based models (Vaswani et al., 2017) have taken centre stage. Multiple works have leveraged Transformer-based pretrained language models such as BERT, GPT and GPT-2 to improve the overall language understanding and generation capabilities of their dialogue agents (Gu et al., 2021; Zhang et al., 2020; Zhao et al., 2019). In addition to latent variable models, different learning approaches such as reinforcement learning (Saleh et al., 2019; Li et al., 2016b) and adversarial learning (Li et al., 2017a)

have also been applied to this task.

2.2 Random Vector Functional Link Neural Networks

The Random Vector Functional Link (RVFL) neural network (Pao and Takefuji, 1992) is essentially a single-layer feed forward neural network with a direct link between the input and output layer. The optimal weights of an RVFL can be obtained iteratively, or through a closed form solution via regularized least squares or the Moore-Penrose pseudoinverse. Prior work has mathematically proven that the RVFL is an efficient and effective universal approximator (Needell et al., 2020). Over the years, multiple RVFL variants including (but not limited to) the deep RVFL (Shi et al., 2021), ensemble deep RVFL (Shi et al., 2021), sparse Pretrained-RVFL (Zhang et al., 2019) and Rotation Forest-RVFL (Malik et al., 2021) have been proposed. Recently, RVFL neural networks have been applied to a broad range of practical tasks across multiple domains such as remote sensing (Dai et al., 2022), malware classification (Elkabbash et al., 2021), medical image classification (Nayak et al., 2020; Katuwal et al., 2019) and even Covid-19 spread forecasting (Hazarika and Gupta, 2020).

3 Methodology

3.1 Randomized Link (RL) Transformer

In this paper, we propose a Randomized Link (RL) Transformer for open-domain dialogue generation. Our proposed model generates the dialogue response based only on the dialogue context. Similar to the standard Transformer, the RL Transformer consists of an encoder and decoder. The encoder maps an input sequence $X = \{x_0, \dots, x_{J-1}\}$ (J refers to the length of the input sequence) to an intermediate representation $Z = \{z_0, \dots, z_{J-1}\}$. The decoder then accepts Z as input and generates the final output $Y = \{y_0, \dots, y_{K-1}\}$ (K refers to the length of the output sequence) token by token, in an auto-regressive manner.

The proposed RL Transformer leverages linear layers with randomly initialized weights to incorporate stochasticity into the response generation process. Our work is largely inspired by the Random Vector Functional Link (RVFL) neural network (Pao and Takefuji, 1992), a single-layer feed forward randomization-based neural network consisting of a fixed randomized hidden layer and a direct link from the input to the output layer. The

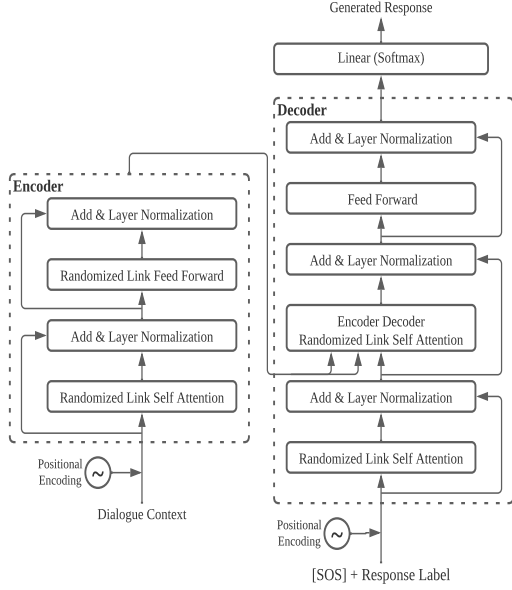


Figure 1: Model architecture of the proposed Randomized Link Transformer.

direct link is implemented by concatenating the input with the randomized hidden layer output, before being fed to the output layer. Only the weights in the output layer are trainable. We incorporate the RVFL architecture into the self-attention and feed forward networks of a Transformer.

Essentially, the Randomized Link Transformer encoder consists of a Randomized Link Self-Attention (RLSA) network and a Randomized Link Feed Forward (RLFF) network. The decoder, on the other hand, consists of a RLSA network, followed by a Encoder-Decoder (ED) RLSA network and a regular Feed Forward (FF) network. Similar to the original Transformer architecture, residual connections and layer normalization is introduced after every RLSA, RLFF, ED RLSA and FF network. An overview is provided in Figure 1.

3.2 Randomized Link Self-Attention

In order to incorporate stochasticity into the response generation process, we introduce a novel Randomized Link Self-Attention (RLSA) network featuring feed forward linear layers with random weights. Similar to the original Transformer architecture, our randomized link self-attention involves mapping a query vector, denoted by Q , and a set of key-value vector pairs, denoted by K and V respectively, to an output. The output is derived by computing the weighted sum of the values, where the weight assigned to each value is computed by taking the dot product of the query with

the corresponding key. For each attention head in the standard transformer architecture, the input is passed to three distinct linear layers, resulting in the query, key and value vectors Q_n , K_n , and V_n , where n refers to an arbitrary attention head.

In RLSA, each input will be fed to a single random linear layer. The sizes of all random linear layers used in the randomized Transformer, denoted by d_{rand} , are identical. Similar to the RVFL, we introduce a direct link by concatenating the output of this layer with the original input. The resultant representation is fed to a separate linear layer with trainable weights. Similarly, for multi-headed RLSA, a distinct Q_n , K_n , and V_n vector is defined for each of the N attention heads. The dimensionality of the Q_n and K_n vectors is denoted by d_k , and the dimensionality of the V_n vector is denoted by d_v . This can be expressed as follows:

$$Q_n = W_{Q_n}([X, W_{Q_n}^r(X)]) \quad (1)$$

$$K_n = W_{K_n}([X, W_{K_n}^r(X)]) \quad (2)$$

$$V_n = W_{V_n}([X, W_{V_n}^r(X)]) \quad (3)$$

where W_Q , W_K , and W_V represent the weights of the trainable linear layers corresponding to the Q , K and V vectors respectively. We use the superscript $()^r$ to denote the randomized matrices, whereby W_Q^r , W_K^r , and W_V^r represent the weights of the randomly initialized linear layers corresponding to the Q_n , K_n and V_n vectors respectively. X denotes the input sequence, and $[\cdot]$ represents the concatenate operation.

The randomized linear layers W_Q^r , W_K^r , and W_V^r are initialized every epoch with Xavier normal initialization (Glorot and Bengio, 2010) i.e., $W_Q^r, W_K^r, W_V^r \sim \mathcal{N}(0, \sigma^2)$ where $\sigma = \gamma \times \sqrt{\frac{2}{d_{hidden} + d_{rand}}}$. The gain value $\gamma = 1.0$. The selection of the standard deviation or the variance of the initialization is vital to model performance. This is because an excessively large variance would result in model divergence, while an excessively small variance would result in a drop in stochasticity. For the randomized layers in RLSA, weights initialized from a normal distribution with a suitable standard deviation would allow the model to converge while maintaining stochasticity. As seen in the equation, for the Xavier normal initialization, the standard deviation of the Gaussian distribution from which the initial weights are sampled is a function of the total number of inputs and outputs. Empirically, we found that the standard deviation

value utilized during Xavier normal initialization would generally outperforms other standard deviation values across all randomized layer sizes.

Then, following the standard Transformer architecture, the dot product of all corresponding Q and K vectors are computed to obtain the attention maps. Then, to attain the score, the softmax function is applied over the dot products divided by the square root of the dimension of the Q and K vectors i.e., $\sqrt{d_k}$. Each of the V vectors is then multiplied with the attained score. This results in the following expression:

$$Z_n = \text{softmax}\left(\frac{Q_n K_n^T}{\sqrt{d_k}}\right) V_n \quad (4)$$

where $\mathbf{T}$ represents the transpose operation. The output of each attention head Z_n is concatenated to form Z :

$$Z = [Z_0, Z_1, Z_2 \cdots Z_{N-1}] \quad (5)$$

where N represents the number of attention heads.

Then, the resulting representation Z is then passed to a single linear layer with randomly initialized weights. Once again, to obtain the encoder output Z , the output of the random layer is concatenated with the original input $\bar{Z}$, and fed to a separate linear layer with trainable weights (direct link).

$$Z = W_Z([X, W_Z^r(Z)]) \quad (6)$$

where W_Z and W_Z^r represent the weights of the trainable linear layer and randomized linear layer used to obtain the encoder output Z respectively. Similarly, the randomized linear layer W_Z^r is initialized every epoch with Xavier normal initialization i.e., $W_Z^r \sim \mathcal{N}(0, \sigma^2)$ where $\sigma = \gamma \times \sqrt{\frac{2}{d_v + d_{rand}}}$. The gain value $\gamma = 1.0$.

For the Encoder Decoder (ED) RLSA in the decoder, the encoder outputs and prior decoder outputs are used as input. The output of the prior decoder layer is used to generate the queries, and the encoder outputs are used to generate the keys and values. An overview of the RLSA network is presented in Figure 2.

3.3 Randomized Link Feed Forward Network

The feed forward network in the standard Transformer consists of a two-layer fully trainable feed forward neural network. Likewise, the Randomized Link Feed Forward (RLFF) network is a two-layer feed forward neural network which features a randomly initialized fixed linear layer with a ReLU

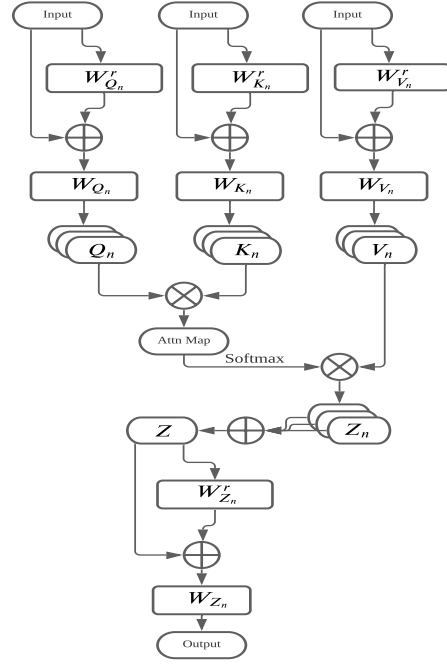


Figure 2: Overview of Randomized Link Self Attention. $\oplus$ refers to the concatenate operation.

activation function followed by a trainable linear layer. Similarly, the direct link is introduced by concatenating the output of the randomized layer with the original inputs, and passing the resultant representation to the trainable layer. In contrast to the RLSA network (Section 3.2), no additional randomized layers are introduced in the RLFF network. The first linear layer is randomized instead. This can be expressed as:

$$RLFF(Z) = W_2([ReLU(W_1^r(Z)), Z]) \quad (7)$$

where W_1^r and W_2 refer to the first random linear layer and the second trainable linear layer respectively. The size of the randomized linear layer W_1^r is represented by d_{ff} , and the size W_2 is d_{hidden} . Unlike the RLSA network, for the RLFF network, the randomized linear layer W_1^r is initialized every epoch with Xavier uniform initialization (Glorot and Bengio, 2010) i.e., $W_Z^r \sim \mathcal{U}(-a, a)$ where $a = \gamma \times \sqrt{\frac{6}{d_{hidden} + d_{ff}}}$. Since the ReLU activation is applied to the layer output, the gain value. The gain value $\gamma = \sqrt{2}$. We found that utilizing a uniform initialization in the RLFF network instead of a normal initialization would result in a slight increase in response diversity. Also, it should be noted that the RLFF network is only utilized in the encoder of the Randomized Transformer. An overview of the RLFF network is presented in Figure 3.

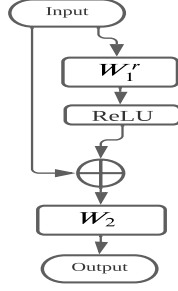


Figure 3: Overview of Randomized Link Feed Forward network. $\oplus$ refers to the concatenate operation.

4 Experiment

4.1 Data

We evaluate the Randomized Link Transformer on the DailyDialogs (Li et al., 2017b) and EmpatheticDialogues (Rashkin et al., 2019) corpora. For the EmpatheticDialogues corpus, the agent is expected to generate an appropriately empathetic response given the dialogue context and emotion label, which is not used in our experiments. The training, validation and test set consists of 19,533, 2,770, and 2,547 dialogues respectively. The DailyDialog corpus consists of general human-written dialogue examples covering a wide range of topics and emotions. Similarly, the provided intent and emotion labels are not used in our experiments. The training, validation and test set consists of 11,118, 1,000 and 1,000 dialogues respectively.

4.2 Implementation

In our experiments, we implement the Randomized Link (RL) Transformer with 4 encoding layers, 4 decoding layers, with 4 attention heads ($N = 4$). Since the 300 dimensional GloVe embedding (Pennington et al., 2014) is used, the hidden dimension $d_{hidden} = 300$. The size of all randomized layers in the RLSA network d_{rand} is fixed at 512. d_k , d_v and d_z were set to 64 for experiments of the DailyDialog corpus, and 256 on the EmpatheticDialogues corpus. For the RLFF network, d_{ff} is set to 2048. Inputs to the RL Transformer consists of the dialogue context (limited to 4 dialogue turns). Responses are generated via greedy decoding. During training, the Adam optimizer (learning rate = 0.00015, batch size = 32) is used.

4.3 Baselines

We implement the following four models in our experiments:

Transformer. We implement a Transformer (Vaswani et al., 2017) with standard self attention and feed forward components. The Transformer parameters are identical to the RL Transformer as described in Section 4.2.

CVAE. Similar to Zhao et al. (2017) and Lin et al. (2020), we implement a Transformer-based CVAE where the latent variable sampled from the latent Gaussian is combined with the output of the encoder before being fed to the decoder. Following Lin et al. (2020), the latent Gaussian is generated by a three-layer MLP with 512 node hidden layers, and the size of the random latent variable is fixed at 300. The remaining Transformer parameters are identical to the RL Transformer as described in Section 4.2.

SVT. We also implement the Sequential Variational Transformer (SVT) proposed in Lin et al. (2020). The SVT replaces the standard Transformer decoder with a variational decoder layer which implicitly generates a distinct latent variable for each position. Similarly, the latent Gaussians are generated via three-layer MLPs with 512 node hidden layers, and the size of the random latent variable is fixed at 300. The remaining Transformer parameters are identical to the RL Transformer described in Section 4.2.

RL Transformer. We implement the proposed RL Transformer with configuration described in Section 3.2.

4.4 Evaluation

4.4.1 Objective Measures

Distinct-N. We use the Distinct-1 and 2 scores, denoted by D-1 and D-2 respectively, to quantify the inter-response diversity of the generated responses. Essentially, the Distinct- n score involves computing the number of distinct n -grams in a given text, and dividing the number by the total number of tokens. The Distinct score is derived based on all generated responses.

MTLD & MATTR. Additionally, we also utilize lexical diversity measures such as the Measure of Textual Lexical Diversity (MTLD) and Moving-Average Type-Token Ratio (MATTR) score. MATTR and MTLD are essentially text length invariant variants of the Token-Type Ratio (TTR). MTLD involves computing the Token-Type Ratio (TTR) for sequentially larger segments of the sentence until a predefined threshold h . On the other hand, deriving MATTR requires averaging

	DailyDialog						
	Length	D-1	D-2	MATTR	MTLD	METEOR	ROUGE-L
Transformer (Vaswani et al., 2017)	6.075	0.004	0.017	0.360	12.461	0.076	0.060
CVAE (Zhao et al., 2017)	11.656	0.035	0.200	0.661	32.927	0.117	0.103
SVT (Lin et al., 2020)	10.244	0.032	0.199	0.580	21.371	0.118	0.108
RL Transformer	7.678	0.050	0.221	0.649	30.049	0.113	0.101
	EmpatheticDialogues						
	Length	D-1	D-2	MATTR	MTLD	METEOR	ROUGE-L
Transformer (Vaswani et al., 2017)	9.757	0.015	0.049	0.396	16.479	0.103	0.116
CVAE (Zhao et al., 2017)	11.367	0.028	0.226	0.728	49.394	0.097	0.084
SVT (Lin et al., 2020)	12.568	0.023	0.240	0.691	36.536	0.105	0.096
RL Transformer	11.808	0.030	0.239	0.734	51.396	0.101	0.085

Table 1: Performance comparison of the proposed RL Transformer to the three baselines on the DailyDialog and EmpatheticDialogues corpora.

the TTR of successive segments of the generated response with a fixed window size w . In this paper, h and w were fixed at 0.72 and 4 respectively. Both MTLD and MATTR are derived based on all generated responses.

ROUGE-L & METEOR. Both ROUGE-L and METEOR compares the generated response to the response label. Computing the ROUGE-L score involves first identifying the Longest Common Subsequence (LCS) between the generated response and response label, followed by computing the harmonic mean of the precision and recall between the LCS and the generated response. METEOR is similarly based on the harmonic mean between precision and recall calculated based on the generated response and response label, with more emphasis placed on recall.

4.4.2 Human Evaluation

For human evaluation, we engage five graduate students (native English speakers) to evaluate the Fluency, Diversity and Coherence of the generated responses. The Fluency criteria measures the naturalness and human-likeness of the generated response. For the Fluency criteria, the evaluators were told to regard the response in isolation, without regard for the dialogue context. The Diversity criteria accounts for the diversity on terms of vocabulary in the generated response. Coherence refers to the contextual coherence of the generated response i.e., the relevance of the generated response in relation to the dialogue context. The evaluators were given the dialogue context, and told to consider the appropriateness of the generated responses with regard to the dialogue context. For each example, they were told to compare a response generated by RL

Transformer and a response generated by either the base Transformer, CVAE or SVT. The superior response in terms of either Fluency, Coherence or Diversity is selected by the evaluator, and either a ‘Win’, ‘Lose’ or ‘Tie’ is assigned to the corresponding model. The percentage of wins, loses or ties for each pair is then computed. Each evaluator was given 50 randomly sampled dialogue contexts and the corresponding responses generated by the implemented models for evaluation.

5 Results and Discussion

5.1 Quantitative Analysis

Automatic metric scores attained on the DailyDialog and EmpatheticDialogues corpora are presented in Table 1. Human evaluation was only conducted on the DailyDialog corpus. Results are presented in Table 3.

Based on the results attained, it can be concluded that the performance of the RL Transformer is comparable to that of latent variable models such as the CVAE and SVT. On the DailyDialog corpus, RL Transformer outperformed all implemented baselines in terms of Distinct-1 and 2 scores. On the EmpatheticDialogues corpus, RL Transformer outperformed all implemented baselines in terms of Distinct-1, MATTR and MTLD.

When it comes to response diversity, the human evaluation results largely corroborate the automatic metric scores. In terms of Diversity, the RL Transformer achieved a high percentage of ‘Wins’ over the standard Transformer, and a relatively high percentage of ‘Ties’ against the latent variable models.

It can also be observed that the RL Transformer also achieved a high percentage of ‘Wins’ over la-

Dialogue Context	User: Hey, Peter , have you had lunch yet?
Transformer	much?
CVAE	I'm sorry , Mom. I was hoping to have my husband.
SVT	No , I didn't . I got to eat it .
RL Transformer	No , I didn't . I just want to go to bed.
Label	No . How about you?
Dialogue Context	User: How did you find your apartment?
Transformer	much?
CVAE	I went there.
SVT	I bought it in the kitchen.
RL Transformer	I had a leaking faucet in the kitchen.
Label	You can check on the bulletin boards at school for local housing.
Dialogue Context	User: Very glad to know something about you , then what are you going to do when you finish. Agent: Oh , I'll go to shanghai to practice there. User: That's a good idea . It must be easy to find a job in shanghai.
Transformer	I'm sorry to see.
CVAE	No , I'm not sure.
SVT	well , that's a great sense. I'll be great.
RL Transformer	I'm not sure I' ll be able to find a job.
Label	I think so , you know there is a great deal of opportunity for business there.

Table 2: Examples of responses generated by the Transformer, CVAE, SVT and RL Transformer models.

tent variable models in terms of Coherence. This is corroborated by the qualitative analysis of the responses provided in the following section. We suspect that low Coherence scores attained by the latent variable models could be partially attributed to the random sampling process. Since the latent Gaussian models the potential dialogue response intents, sampled random variables which deviate significantly from the mean could encompass an irrelevant dialogue intent. When this random variable is fed to the decoder, an incoherent response would likely be generated. Similarly, it can also be observed that the RL Transformer also achieved a high percentage of ‘Wins’ over latent variable models in terms of Fluency. This could be potentially attributed to the random linear layers introduced in the RLSA networks, which serve to improve the overall capability of the RL Transformer.

Additionally, as reported in Liu et al. (2016), we note that the ROUGE and METEOR scores do not correlate with any aspect of human evaluation.

5.2 Qualitative Analysis

We conduct a qualitative analysis by examining the responses generated by each of the implemented

	Fluency			
	Win	Tie	Loss	Kappa
Transformer	23%	63%	14%	0.68
CVAE	71%	19%	10%	0.72
SVT	64%	25%	11%	0.77
	Coherence			
	Win	Tie	Loss	Kappa
Transformer	56%	27%	17%	0.73
CVAE	71%	11%	18%	0.77
SVT	69%	23%	8%	0.67
	Diversity			
	Win	Tie	Loss	Kappa
Transformer	81%	6%	13%	0.62
CVAE	39%	51%	10%	0.59
SVT	46%	39%	15%	0.64

Table 3: Human evaluation results on the DailyDialog corpus. For each criteria (Fluency, Coherence, and Diversity), responses generated by the RL Transformer are compared against responses generated by the Transformer, CVAE and SVT models. The average ‘Win’, ‘Tie’, and ‘Loss’ percentages are presented. Kappa scores largely range from 0.6 to 0.7, indicating substantial to moderate inter-annotator agreement.

	D-1	D-2	MATTR	MTLD
RL Trans	0.050	0.221	0.649	30.049
−RLSS(E)	0.018	0.080	0.523	18.455
−RLFF(E)	0.042	0.173	0.601	24.364
−RLSS(D)	0.020	0.106	0.502	17.237
−ED RLSS	0.026	0.111	0.532	18.231
+RLFF(D)	0.035	0.140	0.577	21.620
RL Trans w/o Links	0.002	0.006	0.220	11.596
RL Trans (Normal)	0.044	0.197	0.634	28.034
RL Trans (Uniform)	0.035	0.185	0.683	28.663

Table 4: Ablation study results. ‘-’ indicates that the corresponding RLSS or RLFF network was replaced with the standard variant. ‘+’ indicates that the standard self-attention or feed forward network was replaced with either RLSS or RLFF.

models. The qualitative analysis largely support observations from our quantitative analysis. The responses generated by the standard transformer were short, generic and repetitive. As expected, responses generated by the latent variable models CVAE and SVT as well as our proposed RL Transformer were noticeably more diverse and less repetitive.

However, numerous responses generated by CVAE and SVT were relatively unnatural due to the relatively poor fluency and contextual coherence. A large number responses generated by the latent variable models had grammatical or structural issues, and a significant number were irrelevant in relation to the dialogue context i.e., out of context. On the other hand, the responses generated by the RL Transformer were significantly more natural and human-like, displaying far fewer grammatical or structural issues and greater relevance with regard to the dialogue context.

Samples of the generated responses along with the corresponding dialogue contexts are provided in Table 2.

5.3 Ablation Study

We also conduct an ablation study, using the DailyDialog corpus, to investigate the contributions of each of the RLSS and RLFF components in the encoder and decoder. Additionally, to examine the importance of the direct links, we implement a variant of the RL Transformer without any direct links. The results of the ablation study are

	D-1	D-2	MATTR	MTLD
64	0.005	0.025	0.426	13.142
128	0.036	0.142	0.527	17.396
256	0.043	0.168	0.628	22.053
512	0.050	0.221	0.649	30.049
1024	0.023	0.100	0.586	19.884

Table 5: Ablation study results for 64, 128, 256, 512, and 1024 nodes in the random layers.

presented in Table 4. In the same table, we also provide the diversity scores for two variants of the RL Transformer where the randomized layers are initialized via Normal initialization ($\mathcal{N}(0.0, 0.1)$) and Uniform initialization ($\mathcal{U}(0.0, 0.01)$) respectively.

Based on the results attained, we can observe that the RLSS network in the encoder, and the RLSS and ED RLSS networks in the decoder have a relatively large impact on response diversity. Substituting the RLSS or ED RLSS networks for the standard self-attention network results in a significant drop on all diversity measures. However, substituting the RLFF in the encoder for a standard feed forward network results in a relatively minor decrease in diversity. Hence, we conclude that stochasticity introduced in the self-attention networks contribute to overall response diversity to a much larger extent compared to the RLFF network. Although, it should also be noted that, substituting the standard feed forward network in the decoder with a RLFF network would result in a slightly lower diversity scores.

Also, the importance of the direct links cannot be overstated. From the results attained by the RL Transformer without links, it is apparent that removal of the direct links would result in extremely low response diversity. Upon closer inspection of the responses generated by this variant of the RL Transformer, we notice that a majority of the generated responses are short, generic and highly repetitive.

In addition, we present the results attained by varying the number of hidden nodes in the randomized layers in Table 5. Intuitively, this implies varying levels of stochasticity. From Table 5, we can observe that increasing the number of neurons in the randomized layer would generally result in an increase in diversity. This can be attributed to an increase in stochasticity due to an increase in the number of randomized weights. However, there

is a significant drop in diversity when the size of the randomized layers exceed 512. In this case, the model fails to learn effectively, and generates meaningless, generic, and repetitive responses instead.

6 Conclusion

In this paper, we have proposed a novel RL Transformer which successfully improves response diversity in the task of open-domain dialogue generation. This is achieved by inducing stochasticity in the self-attention and the feed forward networks of a Transformer via randomized layers and direct links. Experimental results on the DailyDialog and EmpatheticDialogues corpora show that, compared to latent variable models, the RL Transformer achieved comparable levels of diversification while further improving on contextual coherence and fluency. In the future, the RL Transformer can be adapted for related dialogue generation tasks such as personalized, empathetic or topical dialogue generation.

References

- Kris Cao and Stephen Clark. 2017. [Latent variable dialogue models and their diversity](#). In *Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics: Volume 2, Short Papers*, pages 182–187, Valencia, Spain. Association for Computational Linguistics.
- Yu Cao, Liang Ding, Zhiliang Tian, and Meng Fang. 2021. [Towards efficiently diversifying dialogue generation via embedding augmentation](#). In *ICASSP 2021 - 2021 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 7443–7447.
- Qijun Dai, Gong Zhang, Zheng Fang, and Biao Xue. 2022. [Sar target recognition with modified convolutional random vector functional link network](#). *IEEE Geoscience and Remote Sensing Letters*, 19:1–5.
- Emad T. Elkabbash, Reham R. Mostafa, and Sherif I. Barakat. 2021. [Android malware classification based on random vector functional link and artificial jellyfish search optimizer](#). *PLOS ONE*, 16(11):1–22.
- Jun Gao, Wei Bi, Xiaojiang Liu, Junhui Li, Guodong Zhou, and Shuming Shi. 2019. [A discrete CVAE for response generation on short-text conversation](#). In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pages 1898–1908, Hong Kong, China. Association for Computational Linguistics.
- Xavier Glorot and Yoshua Bengio. 2010. [Understanding the difficulty of training deep feedforward neural networks](#). In *Proceedings of the Thirteenth International Conference on Artificial Intelligence and Statistics*, volume 9 of *Proceedings of Machine Learning Research*, pages 249–256, Chia Laguna Resort, Sardinia, Italy. PMLR.
- Xiaodong Gu, Kang Min Yoo, and Jung-Woo Ha. 2021. [DialogBERT: Discourse-aware response generation via learning to recover and rank utterances](#).
- Barenaya Bikash Hazarika and Deepak Gupta. 2020. [Modelling and forecasting of covid-19 spread using wavelet-coupled random vector functional link networks](#). *Applied Soft Computing*, 96:106626.
- Rakesh Katuwal, P. N. Suganthan, and M. Tanveer. 2019. [Random vector functional link neural network based ensemble deep learning](#).
- Jing Yang Lee., Kong Aik Lee., and Woon Seng Gan. 2022. [Dlvgen: A dual latent variable approach to personalized dialogue generation](#). In *Proceedings of the 14th International Conference on Agents and Artificial Intelligence - Volume 2: ICAART*, pages 193–202. INSTICC, SciTePress.
- Jiwei Li, Michel Galley, Chris Brockett, Jianfeng Gao, and Bill Dolan. 2016a. [A diversity-promoting objective function for neural conversation models](#). In *Proceedings of the 2016 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 110–119, San Diego, California. Association for Computational Linguistics.
- Jiwei Li, Will Monroe, Alan Ritter, Dan Jurafsky, Michel Galley, and Jianfeng Gao. 2016b. [Deep reinforcement learning for dialogue generation](#). In *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing*, pages 1192–1202, Austin, Texas. Association for Computational Linguistics.
- Jiwei Li, Will Monroe, Tianlin Shi, Sébastien Jean, Alan Ritter, and Dan Jurafsky. 2017a. [Adversarial learning for neural dialogue generation](#). In *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*, pages 2157–2169, Copenhagen, Denmark. Association for Computational Linguistics.
- Mei Li, Jiajun Zhang, Xiang Lu, and Chengqing Zong. 2021. [Dual-view conditional variational auto-encoder for emotional dialogue generation](#). *ACM Trans. Asian Low-Resour. Lang. Inf. Process.*, 21(3).
- Qintong Li, Hongshen Chen, Zhaochun Ren, Pengjie Ren, Zhaopeng Tu, and Zhumin Chen. 2020a. [EmpDG: Multi-resolution interactive empathetic dialogue generation](#). In *COLING*, pages 4454–4466, Barcelona, Spain (Online). International Committee on Computational Linguistics.

- Shifeng Li, Shi Feng, Daling Wang, Kaisong Song, Yifei Zhang, and Weichao Wang. 2020b. Emoelicitator: An open domain response generation model with user emotional reaction awareness. In *Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence, IJCAI-20*, pages 3637–3643. International Joint Conferences on Artificial Intelligence Organization. Main track.
- Yanran Li, Hui Su, Xiaoyu Shen, Wenjie Li, Ziqiang Cao, and Shuzi Niu. 2017b. [DailyDialog: A manually labelled multi-turn dialogue dataset](#). In *Proceedings of the Eighth International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 986–995, Taipei, Taiwan. Asian Federation of Natural Language Processing.
- Zhaojiang Lin, Genta Indra Winata, Peng Xu, Zihan Liu, and Pascale Fung. 2020. Variational transformers for diverse response generation. *arXiv preprint arXiv:2003.12738*.
- Chia-Wei Liu, Ryan Lowe, Iulian Serban, Mike Noseworthy, Laurent Charlin, and Joelle Pineau. 2016. [How NOT to evaluate your dialogue system: An empirical study of unsupervised evaluation metrics for dialogue response generation](#). In *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing*, pages 2122–2132, Austin, Texas. Association for Computational Linguistics.
- Ashwani Kumar Malik, M.A. Ganaie, M. Tanveer, and P.N. Suganthan. 2021. [A novel ensemble method of rvfl for classification problem](#). In *2021 International Joint Conference on Neural Networks (IJCNN)*, pages 1–8.
- Ryo Nakamura, Katsuhito Sudoh, Koichiro Yoshino, and Satoshi Nakamura. 2018. [Another diversity-promoting objective function for neural dialogue generation](#). *CoRR*, abs/1811.08100.
- Deepak Ranjan Nayak, Ratnakar Dash, Banshidhar Majhi, Ram Bilas Pachori, and Yudong Zhang. 2020. [A deep stacked random vector functional link network autoencoder for diagnosis of brain abnormalities and breast cancer](#). *Biomedical Signal Processing and Control*, 58:101860.
- Deanna Needell, Aaron A. Nelson, Rayan Saab, and Palina Salanevich. 2020. [Random vector functional link networks for function approximation on manifolds](#). *CoRR*, abs/2007.15776.
- Y.-H. Pao and Y. Takefuji. 1992. [Functional-link net computing: theory, system architecture, and functionalities](#). *Computer*, 25(5):76–79.
- Jeffrey Pennington, Richard Socher, and Christopher Manning. 2014. [GloVe: Global vectors for word representation](#). In *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 1532–1543, Doha, Qatar. Association for Computational Linguistics.
- Hannah Rashkin, Eric Michael Smith, Margaret Li, and Y-Lan Boureau. 2019. [Towards empathetic open-domain conversation models: A new benchmark and dataset](#). In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 5370–5381, Florence, Italy. Association for Computational Linguistics.
- Abdelrhman Saleh, Natasha Jaques, Asma Ghandeharion, Judy Shen, and Rosalind Picard. 2019. Hierarchical reinforcement learning for open-domain dialog. *arXiv preprint arXiv:1909.07547*.
- Lifeng Shang, Zhengdong Lu, and Hang Li. 2015. [Neural responding machine for short-text conversation](#). In *Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 1577–1586, Beijing, China. Association for Computational Linguistics.
- Qiushi Shi, Rakesh Katuwal, P.N. Suganthan, and M. Tanveer. 2021. [Random vector functional link neural network based ensemble deep learning](#). *Pattern Recognition*, 117:107978.
- Kihyuk Sohn, Honglak Lee, and Xinchun Yan. 2015. [Learning structured output representation using deep conditional generative models](#). In *Advances in Neural Information Processing Systems*, volume 28. Curran Associates, Inc.
- Haoyu Song, Weinan Zhang, Yiming Cui, Dong Wang, and Ting Liu. 2019. Exploiting persona information for diverse generation of conversational responses. In *IJCAI*.
- Alessandro Sordani, Michel Galley, Michael Auli, Chris Brockett, Yangfeng Ji, Margaret Mitchell, Jian-Yun Nie, Jianfeng Gao, and Bill Dolan. 2015. [A neural network approach to context-sensitive generation of conversational responses](#). In *Proceedings of the 2015 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 196–205, Denver, Colorado. Association for Computational Linguistics.
- Ponnuthurai N. Suganthan and Rakesh Katuwal. 2021. [On the origins of randomization-based feedforward neural networks](#). *Applied Soft Computing*, 105:107239.
- Ilya Sutskever, Oriol Vinyals, and Quoc V. Le. 2014. Sequence to sequence learning with neural networks. In *Proceedings of the 27th International Conference on Neural Information Processing Systems - Volume 2*, NIPS’14, page 3104–3112, Cambridge, MA, USA. MIT Press.
- Ayaka Ueyama and Yoshinobu Kano. 2020. [Diverse dialogue generation with context dependent dynamic loss function](#). In *Proceedings of the 28th International Conference on Computational Linguistics*, pages 4123–4127, Barcelona, Spain (Online). International Committee on Computational Linguistics.

- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. [Attention is all you need](#). In *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc.
- Yiru Wang, Pengda Si, Zeyang Lei, and Yujiu Yang. 2020. [Topic enhanced controllable cvae for dialogue generation \(student abstract\)](#). *Proceedings of the AAAI Conference on Artificial Intelligence*, 34(10):13955–13956.
- Bowen Wu, MengYuan Li, Zongsheng Wang, Yifu Chen, Derek F. Wong, Qihang Feng, Junhong Huang, and Baoxun Wang. 2020. [Guiding variational response generator to exploit persona](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 53–65, Online. Association for Computational Linguistics.
- Haiqin Yang, Xiaoyuan Yao, Yiqun Duan, Jianping Shen, Jie Zhong, and Kun Zhang. 2021. Progressive open-domain response generation with multiple controllable attributes. In *Proceedings of the Thirtieth International Joint Conference on Artificial Intelligence, IJCAI-21*, pages 3279–3285. International Joint Conferences on Artificial Intelligence Organization. Main Track.
- Yizhe Zhang, Siqi Sun, Michel Galley, Yen-Chun Chen, Chris Brockett, Xiang Gao, Jianfeng Gao, Jingjing Liu, and Bill Dolan. 2020. [DIALOGPT : Large-scale generative pre-training for conversational response generation](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics: System Demonstrations*, pages 270–278, Online. Association for Computational Linguistics.
- Yongshan Zhang, Jia Wu, Zhihua Cai, Bo Du, and Philip S. Yu. 2019. [An unsupervised parameter learning model for rvfl neural network](#). *Neural Networks*, 112:85–97.
- Tiancheng Zhao, Ran Zhao, and Maxine Eskenazi. 2017. [Learning discourse-level diversity for neural dialog models using conditional variational autoencoders](#). In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 654–664, Vancouver, Canada. Association for Computational Linguistics.
- Xue Zhao, Ying Zhang, Wenya Guo, and Xiaojie Yuan. 2019. [Bert for open-domain conversation modeling](#). In *2019 IEEE 5th International Conference on Computer and Communications (ICCC)*, pages 1532–1536.
- Xianda Zhou and William Yang Wang. 2018. [MojiTalk: Generating emotional responses at scale](#). In *ACL*, pages 1128–1137, Melbourne, Australia. Association for Computational Linguistics.