

SS-MPC: A Sequence-Structured Multi-Party Conversation System

Anonymous ACL submission

Abstract

Recent Multi-Party Conversation (MPC) models typically rely on graph-based approaches to capture dialogue structures. However, these methods have limitations, such as information loss during the projection of utterances into structural embeddings and constraints in leveraging pre-trained language models directly. In this paper, we propose **SS-MPC**, a response generation model for MPC that eliminates the need for explicit graph structures. Unlike existing models that depend on graphs to analyze conversation structures, SS-MPC internally encodes the dialogue structure as a sequential input, enabling direct utilization of pre-trained language models. Experimental results show that **SS-MPC** achieves **15.60% BLEU-1** and **12.44% ROUGE-L** score, outperforming the current state-of-the-art MPC response generation model by **3.91%p** in **BLEU-1** and **0.62%p** in **ROUGE-L**. Additionally, human evaluation confirms that SS-MPC generates more fluent and accurate responses compared to existing MPC models.

1 Introduction

The rapid development of the Internet and the social media platforms have changed the way people communicate with each other and created new forms of interaction. In particular, Multi-Party Conversation (MPC), conversation in which multiple people freely exchanges opinions at the same time, is increasingly becoming common. MPC occurs on many platforms, such as group chats, online forums, and comment sections on social media. Recent research trends show that analysis and response generation on MPC is in its infancy, and the importance and need for them is increasingly being emphasized (Park and Lim, 2020; Anjum et al., 2020).

MPC have flexibility to allow multiple speakers to participate in a conversation in no particular order or rule. These attributes of MPC complicate the

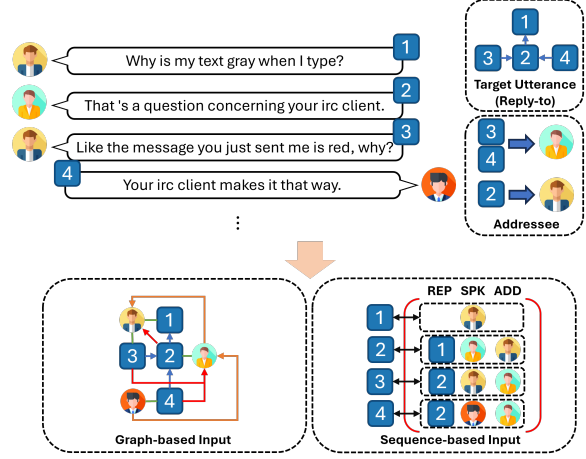


Figure 1: An example of data in MPC Dataset (Ubuntu IRC Benchmark Dataset). The dataset is constructed of context and structural information. Context consists of utterances, and structural information consists of speaker information, target-utterance relation and addressee relation of each utterance.

flow and structure of the conversation compared to one-on-one conversations, and create additional challenges in understanding the context and intent of the utterances and speakers. Each speaker brings a unique context and intent to the conversation, and he or she must understand and process the context in which a particular utterance is being delivered to whom. However, because the addressee of an utterance is often unclear in MPC, predicting or analyzing the structure of the dialogue is one of the major challenges for MPC. In addition, because multiple participants are simultaneously expressing their opinions and interacting, the topic and flow of the conversation are likely to change frequently. Unlike one-on-one conversations, these characteristics create the additional challenge of closely tracking and managing the context of each utterance in MPC. For these reasons, developing systems for generating responses to MPC is considered one of the most challenging areas of current

dialogue system research.

To address this complexity, MPC datasets typically include more structural conversation information for each utterance (Lowe et al., 2015). For example, as shown in Figure 1, each utterance contains the structural information such as speaker information, which tells us who is speaking, the addressee information, which tells us to whom the utterance is addressed, and the target-utterance information, which indicates which utterance the current utterance is responding to. The target-utterance relationship is typically linked to only one previous utterance in the conversation history, and the addressee is semantically the same as the speaker of the target-utterance. This structural information is critical for dialogue systems to understand and learn about the complexity of MPC, and it helps MPC response generation models generate responses that are appropriate for a given context.

However, traditional language models using sequential input have found that it is very difficult to express the structure of these conversations. Recent work (Gu et al., 2022, 2023b) has attempted to solve this problem using graphs. By representing utterances, speakers, and the relationships between utterances and speakers as graphs, and interpreting them through a graph encoder, it was possible to train a model to recognize the structure of a conversation and generate responses accordingly. But this still limits the use of pre-trained language models itself, since there are no such pre-trained graph-encoder models tuned properly to analyze the structure of conversations. If we partially employ randomly initialized graph-encoder, it may break the embedding space of the entire pre-trained model.

In this paper, we introduce Sequence-Structured MPC (SS-MPC), a response generation system for multi-party conversations that leverages the encoder-decoder architecture of the transformer while replacing the role of a graph encoder with a soft prompt. Instead of explicitly encoding conversational structures using a graph encoder, SS-MPC effectively integrates structural information through well-designed soft prompts within the standard encoder-decoder framework. This approach eliminates the need for additional model components, accelerates the training process, and achieves superior performance compared to existing models.

Futhermore, unlike additinal models that require MPC analysis for the response generation, SS-MPC has the advantage of being able to ana-

lyze conversations and generate responses at once by end-to-end, enabling immediate usage in real-world MPC environments. By learning the conversation structure such as the correct target-utterance and addressee information for each utterance during the training process, the model can make appropriate inferences and generate the final response even when some of the correct target-utterance and addressee information is omitted during the actual inference process. This can be done because the SS-MPC contains the post-training process with the way masking the information partially, which means that the model itself is already trained for the ability to predict the omitted target-utterance or addressee information.

Our contributions are summarized as follows:

- We propose a novel method to train language models for MPC response generation without graph structures, which leverages sequence-structured inputs to internally represent the interaction flow in the dialogue.
- The proposed model can be used in real-world MPC environments easily because the model can simultaneously analyze conversations and generate responses using an end-to-end framework.
- Experimental results show that the proposed model performs better than the previous SOTA model. In addition, our various analysis results provide directions for future research in multi-party dialogues.

2 Related Work

2.1 Multi-Party Dialogue Structural Analysis

The task of predicting relationships between speakers to analyze the structure of MPC began in 2016. Ouchi and Tsuboi (2016) first proposed the addressee prediction and utterance selection task. To study this task, they first hand-created a dataset of MPC using log transcripts from Ubuntu IRC channels, and then utilized RNNs to perform the proposed task. Later, Zhang et al. (2018) proposed SI-RNN, which updates speaker embeddings based on roles for addressee prediction. Meng et al. (2017) also proposed a speaker classification task to model the relationships between speakers. Meanwhile, for the MPC response selection task, Wang et al. (2020) proposed to track dynamic topics, and then a who-to-whom (W2W) model (Le et al., 2019) was

proposed to predict the addressees of all utterances in a conversation. Gu et al. (2021) proposed the MPC-BERT model, which utilizes multiple MPC learning methods to learn the complex interactions between recent utterances and interlocutors, and it performs post-training for MPC tasks. In addition, Gu et al. (2023a) proposed the GIFT model to help fine-tuning for MPC tasks with only simple scalar parameters on the attentions.

However, all the methodologies proposed in the above works have the limitation that they utilize the utterances to predict the addressee information of each utterance. This makes it difficult to utilize these methodologies for response generation models in real-world MPC environments.

2.2 Multi-Party Dialogue Response Generation

Along with these MPC tasks, there has been a parallel research on MPC response generation, which is the task of generating responses to a multi-party dialog. Hu et al. (2019) proposed a graph structure network (GSN) to model the graphical information flow for response generation. Later, Heter-MPC (Gu et al., 2022) was proposed to model complex interactions between utterances and interlocutors as graphs. This paper used graphs with two types of nodes and six types of edges to model the structure of multi-party conversations. Li and Zhao (2023) utilized the Expectation-Maximization (EM) algorithm in pre-training to predict the missing addressee information in the dataset. However, they still suffer from the drawback that the fine-tuning process only allows for an ideal setup where all addressees are labeled. To overcome this, MADnet (Gu et al., 2023b) utilizes the EM algorithm in the model of Gu et al. (2022) to directly predict and supplement the missing addressee information for training and response generation.

3 Methodology

In this paper, we propose SS-MPC, a novel MPC response generation model with the encoder-decoder structure of transformer. Here, we describe its input, output, and the training process.

3.1 Preliminaries

In a typical response generation task, the goal is to generate a final response $\bar{r}$ for a given conversation history h . In addition, the traditional MPC response generation task utilizes not only the dialogue history but also the dialogue structural information

C . C consists of the speaker c_i , addressee a_i , and target-utterance u_i for each of the utterances.

$$C = \{c_1, c_2, \dots, c_n\} \quad (1)$$

$$c_i = \{s_i, a_i, u_i\} \quad (2)$$

where the number of the utterances is n and $1 \leq i \leq n$.

In general, the MPC datasets provide C , and the MPC models perform the task of finding the most appropriate response $\bar{r}$ based on the given h and C information. The response tokens are generated in an auto-regressive way. This can be formulated as follows:

$$\bar{r} = \underset{r}{\operatorname{argmax}} P(r|h, C; \theta)$$

$$= \underset{r}{\operatorname{argmax}} \prod_{t=1}^{|r|} P(r_t|h, C, r_{<t}; \theta) \quad (3)$$

where r_t means the t -th token, and $r_{<t}$ means the previous tokens of the final response r .

The existing MPC models have utilized graph-based models to use C . In this process, there can be a loss of information when aligning the embedding space of h with C and using a graph encoder after random initialization.

$$\bar{r} = \underset{r}{\operatorname{argmax}} P_{Dec}(r|GraphEnc(h, C); \theta) \quad (4)$$

where P_{Dec} means the probability of each token computed by the decoder of the model, and $GraphEnc(\cdot)$ means the embedding created by the graph encoder of the model.

In the case of SS-MPC, we use the sequence-structured input of each utterance, so we can utilize the full information of dialogue history and dialogue structural information without using a graph encoder.

$$\bar{r} = \underset{r}{\operatorname{argmax}} P_{Dec}(r|Enc(h, s(C)); \theta) \quad (5)$$

where $s(\cdot)$ is the dialogue structuralization, which means transforming original input as the sequence-structured input containing the dialogue structure information internally.

Furthermore, we should consider the situation where the lack of structural information exists in MPC. Graph-encoder needs a fully connected graph since the lack of edges means that unconnected nodes can confuse the model in encoding. But SS-MPC does not need any change in model

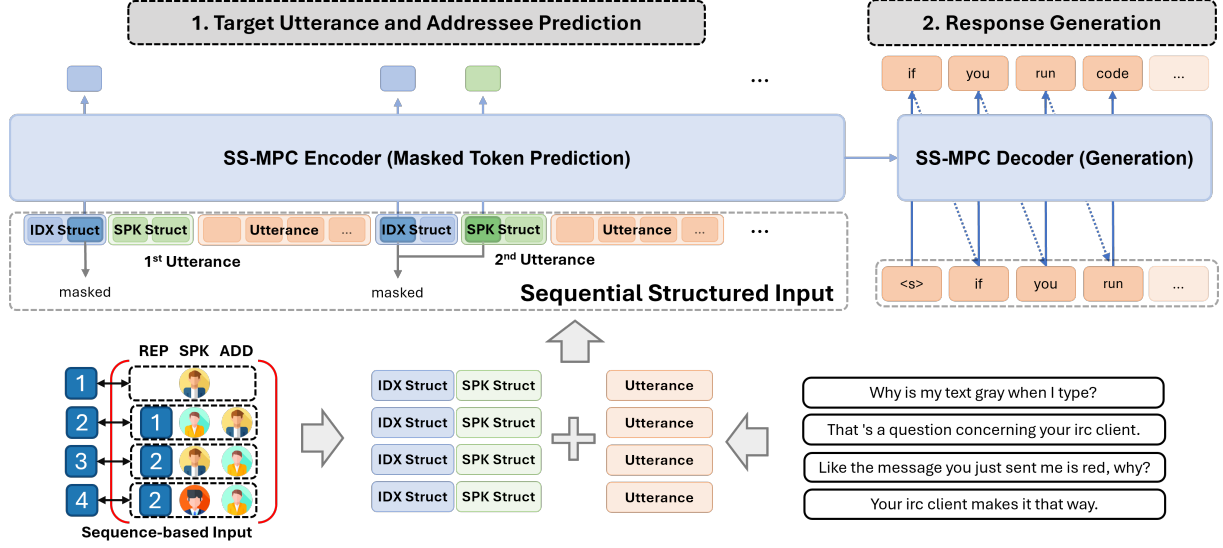


Figure 2: The overview of the SS-MPC. The encoder part is expected to analyze the dialogue and predict the structural information in dialogue. The decoder part is expected to generate the final response with using the information analyzed in encoder.

structure or additional training in this situation because the model is already trained with the way masking the information partially. The inference process of the SS-MPC is as follows:

$$\bar{r}, \bar{C} = \underset{r}{\operatorname{argmax}} P(r|h, s(C_{omit}); \theta) \quad (6)$$

where C_{omit} is the partially omitted structural information in MPC, and $\bar{C}$ means the predicted structural information for C_{omit} .

3.2 Overview of SS-MPC

Figure 2 shows the overview of the SS-MPC. It utilizes the encoder-decoder structure of the transformer because the transformer encoder is commonly used to analyze the structure of conversations (Shen et al., 2020; Mehri et al., 2019), and we want to leverage the strengths of these encoders to analyze the structure of conversations while allowing the decoder to focus on generating the actual responses.

The first step to utilize sequence-structured input instead of graph-structured input is to insert the conversation structure into a sequential form. To do this, we add three soft prompt tokens to the model tokenizer as MPC structure tokens to represent the structure of the conversation by adding information about each utterance. The embeddings of new tokens act as soft prompts for each utterance, which reflect the structural information of the conversation.

Then we post-train only the encoder with the task of predicting dialogue structure to improve the encoder’s ability to interpret the meaning of the added structure tokens and analyze the context with the dialogue structure. Through this process, the model not only learns the meaning of the added MPC structure tokens, but also learns to predict the partially omitted dialogue structure information of each utterance using the information of previous utterances and speakers.

3.3 MPC Structure tokens

To represent the speaker, addressee, and target-utterance for each utterance as a dialogue structure, three kinds of soft prompt tokens are added to the model tokenizer, called by MPC structure tokens; their embeddings are randomly initialized before training. The added MPC structure tokens are as follows:

- Index structure token
- Speaker structure token
- Structure masking token

Index structure token These tokens are developed to distinguish the order of the conversation; they are designed by indicating the order of the each utterance by number, such as "[IDX₁]", "[IDX₂]", ..., "[IDX_n]". The target-utterance can be also represented by specifying the index of the target-utterance.

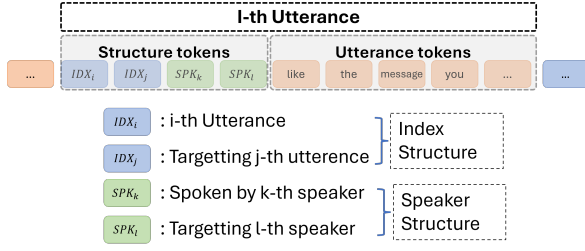


Figure 3: An example of the sequence-structure template for an utterance. Two index structure tokens which represents the utterance’s index and the target-utterance’s index, and two speaker structure tokens which represents speaker and addressee of the utterance are added as prefix to the tokenized utterance tokens.

Speaker structure token These tokens are to distinguish speakers in a conversation based on the order in which they appear; they are designed to distinguish between speakers, such as "[SPK₁]", "[SPK₂]", ..., "[SPK_m]". The addressee information for each utterance can be expressed by specifying the speaker information via the corresponding token.

Structure masking token "[Mask_{IDX}]" and "[Mask_{SPK}]" tokens are added to mask MPC structure information. The "[Mask_{IDX}]" token masks the index and the target-utterance information, and the "[Mask_{SPK}]" token masks the speaker and the addressee information in a given utterance.

Since there is not a target-utterance and addressee for the first utterance, we add "[IDX_{None}]" token and "[SPK_{None}]" token to indicate that it has no target-utterance or addressee information.

3.4 Dialogue Structuralization

Dialogues are sequence-structured using MPC structure tokens that are added to express the structure of the conversation in a sequential form. The entire conversation can be broken down into utterances, and each of the utterances consists of the structure tokens and utterance tokens. Note that, for the final response, the utterance tokens are omitted because it is the answer which the model should generate.

In Figure 3, an example is shown for the i -th utterance of the entire conversation. The token that indicates that the i -th utterance is labeled by "[IDX_i]". This i -th utterance is responding to the j -th utterance, which labeled by the second token "[IDX_j]". Furthermore, the i -th utterance is being uttered by the speaker- k and is being answered

to the speaker- l , which is represented by the tokens "[SPK_k]" and "[SPK_l]" in the following sequence. Thus we can set the sequence-structure template of i -th utterance as follows:

$$S_i = (\overbrace{[IDX_i]; [IDX_j]; [SPK_k]; [SPK_l]}^{\text{structure tokens}}; \underbrace{[token_1]; [token_2] \dots}_{\text{utterance tokens}}) \quad (7)$$

where $1 \leq i \leq n$ for dialogue with n utterances.

Then the model should generate the final response, so its structure information is inputted. The sequence-structure template is set as follows:

$$S_r = (\overbrace{[IDX_r]; [IDX_t]; [SPK_s]; [SPK_a]}^{\text{response structure tokens}}) \quad (8)$$

where "[IDX_r]", "[IDX_t]", "[SPK_s]", and "[SPK_a]" tokens means the index, target-utterance index, speaker, and addressee of the current response utterance.

The sequence-structure templates of whole utterances in a dialogue are concatenated as a sequence-structured input $S = (S_1; S_2; \dots; S_{n-1}; S_n; S_r)$. In addition, the structural information of the input can be masked using the "[Mask_{IDX}]" and "[Mask_{SPK}]" tokens. This masking approach is performed when target-utterance and addressee information has to be predicted in the conversation. In particular, the encoder is post-trained using the masking approach.

3.5 Post-training for Encoder

SS-MPC with the addition of soft prompt tokens needs to carry out post-training to obtain better embeddings of soft prompt tokens containing semantic and contextual information from context tokens. In particular, masking is performed with a probability of hyper-parameter $p\%$ for the structure tokens of each utterances, and the encoder predicts the correct answer for the masked structure tokens. To predict the structure token, the LM head of post-training for encoder shares the parameter with the LM head in decoder since they generate the same type of tokens; For the utterance tokens, we train the encoder to generate the tokens itself. The formulation of the loss function for post-training is as follows:

$$\mathcal{L}_{post} = - \sum_i \log P(x_i | X_{masked}; \phi) \quad (9)$$

where i means each position of the input, x_i denotes the original encoder input token of the i -th position, for the masked encoder input X_{masked} . And the ϕ is the parameter of the SS-MPC encoder.

3.6 Fine-Tuning Model

SS-MPC is fine-tuned after post-training to perform the task of generating the final response. Fine-tuning is same as the learning process of a typical transformer encoder-decoder model. The difference is that the SS-MPC utilizes sequence-structured input.

$$\mathcal{L} = -\sum_{i=1}^n \log P(r_i | r_{<i}, X; \theta) \quad (10)$$

where r_i is the i -th final response token, X is the sequence-structured encoder input, and θ is the parameter of the SS-MPC.

4 Experiments

Dataset To evaluate the performance of the proposed SS-MPC model, we utilize the Ubuntu IRC benchmark dataset, which has originally released by Ouchi and Tsuboi (2016) and Hu et al. (2019), and has been widely using for various MPC tasks. This dataset comprises user conversations from the Internet Relay Chat (IRC) channel of the Ubuntu homepage.

Two Ubuntu IRC Benchmark datasets are used in the experiments as follows:

Ubuntu IRC (2016): The dataset released by Ouchi and Tsuboi (2016)¹ has some missing structural information in the dataset. We construct the sequence-structured input with masking those missing information. This dataset is categorized into three subsets based on session length (Len-5, Len-10, and Len-15). We employ the Len-5 subset, following the settings of previous studies.

Ubuntu IRC (2019): The dataset released by Hu et al. (2019)² includes all structure information for every utterances. The dataset is used for post-training SS-MPC both in Ubuntu IRC (2016) and Ubuntu IRC (2019). Further details on the datasets can be found in the Appendix A.

¹We adopt the refined version provided by Le et al. (2019), which is released on <https://github.com/lxchtan/HeterMPC> (Gu et al., 2022)

²We adopt the re-emblemnted processed version of (Gu et al., 2022), which is released on <https://github.com/lxchtan/HeterMPC>

Evaluation Metrics To evaluate SS-MPC, we just follow previous research and measure its performance using BLEU-1 through BLEU-4, METEOR, and ROUGE-L score for the final response. All metrics are computed using the Hugging Face evaluate library³ (Wolf et al., 2020).

Baselines For the backbone model, we use BART (Lewis et al., 2020) as a widely recognized transformer-based encoder-decoder model. BART leverages the encoder-decoder architecture that are well-suited for response generation as well as other generative tasks such as summarization and machine translation.

We compare our approach against the following models: (1) GSN (Hu et al., 2019). GSN’s core architecture consists of an utterance-level graph-structured encoder. (2) GPT-2 (Radford et al., 2019), a unidirectional pre-trained language model. Following its original setup, all context utterances and response are concatenated by using a special token "[SEP]" as input. We also compare ConvMPC with the (3) HeterMPC (Gu et al., 2022) and (4) MADnet (Gu et al., 2023b), which are known as SOTA models among the current MPC response generation models. The HeterMPC model structures the speaker, target utterance, and addressee relationships in the form of a heterogeneous graph to model complex MPC. To analyze the structured data, it utilizes a heterogeneous graph encoder structure that utilizes Graph Attention (GAT) operations. In the case of the MADnet model, it is a model that slightly modifies the graph input of the existing HeterMPC model and adds the EM-algorithm methodology to generate a response by inferring the missing data from the existing data through the EM-algorithm.

We further evaluate our approach on decoder based Large Language Model (LLM) in appendix C.

Implementation Details Model parameters are initialized with BART-large (Lewis et al., 2020), which were implemented in Hugging Face’s Transformers library⁴ (Wolf et al., 2020). We use AdamW (Loshchilov and Hutter, 2017) for optimization, with an initial learning rate of $1e-5$ that decayed linearly. The model is trained on the two Ubuntu IRC benchmark training sets, with a maximum of 10 epochs for post-training and fine-tuning

³<https://huggingface.co/docs/evaluate/index>

⁴<https://huggingface.co/facebook/bart-large>

Dataset	Models	BLEU-1	BLEU-2	BLEU-3	BLEU-4	METEOR	ROUGE-L
Ubuntu IRC (2016)	GSN [†] (Hu et al., 2019)	10.23	3.57	1.70	0.97	4.10	9.91
	GPT-2 (Radford et al., 2019)	8.86	2.69	1.11	0.61	7.40	8.53
	BART (Lewis et al., 2020)	11.76	4.86	2.97	2.21	8.91	9.86
	MADNet (Gu et al., 2023b)	11.82	4.58	2.65	1.91	9.78	10.61
	SS-MPC	13.40	5.87	3.60	2.65	10.15	11.14
Ubuntu IRC (2019)	GSN [†] (Hu et al., 2019)	6.32	2.28	1.10	0.61	3.27	7.39
	GPT-2 (Radford et al., 2019)	10.85	3.76	1.61	0.84	9.00	7.24
	BART (Lewis et al., 2020)	12.71	4.52	2.13	1.25	8.75	10.01
	HeterMPC (Gu et al., 2022)	10.29	3.68	1.71	0.96	8.79	11.22
	MADNet (Gu et al., 2023b)	11.69	4.57	2.33	1.45	9.48	11.82
	SS-MPC (Ours)	15.60	6.62	3.67	2.44	10.92	12.44

Table 1: Performance comparison of different models on two Ubuntu IRC Benchmark datasets (Ouchi and Tsuboi, 2016; Hu et al., 2019) with various metrics. The performance result of GSN[†] is cited from Gu et al. (2023b).

Human Evaluation	Score	Kappa
Gold Label	1.91	0.51
GPT-2 (Hu et al., 2019)	0.6	0.50
BART (Lewis et al., 2020)	1.50	0.48
MADNet (Gu et al., 2023b)	1.57	0.46
SS-MPC (Ours)	1.84	0.55

Table 2: Human Evaluation results on Ubuntu IRC Benchmark test set of Ubuntu IRC (2019) Hu et al. (2019) with several MPC response generation models.

individually. We use a batch size of 8 with 2 gradient accumulation steps and select the best model based on validation performance for testing. The maximum length of the generated output is set to 50 tokens just following previous studies.

4.1 Main Results

Table 1 shows model performances on two Ubuntu IRC Benchmark datasets, **Ubuntu IRC (2016)** (Ouchi and Tsuboi, 2016) and **Ubuntu IRC (2019)** (Hu et al., 2019). SS-MPC achieves a significant performance improvement on both datasets compared to previous models. On the **Ubuntu IRC (2016)** dataset, SS-MPC outperforms the previous state-of-the-art (SOTA) model, MADNet, by 1.58%p in BLEU-1, 1.29%p in BLEU-2, 0.95%p in BLEU-3, 0.74%p in BLEU-4, 0.37%p in METEOR, and 0.53%p in ROUGE-L.

Similarly, on the **Ubuntu IRC (2019)** dataset, SS-MPC also surpasses MADNet by 3.91%p in BLEU-1, 2.05%p in BLEU-2, 1.34%p in BLEU-3, 0.99%p in BLEU-4, 1.44%p in METEOR, and 0.62%p in ROUGE-L. You can see the response examples generated by each model in Appendix B.

4.2 Human Evaluation

Since quantitative metrics alone may not fully capture the quality of generated responses, we also conduct a human evaluation. Specifically, we ran-

Masking Rates	BLEU-1	METEOR	ROUGE-L
w/o post-training	14.22	9.87	11.84
p=25% (SS-MPC)	15.60	10.92	12.44
p=50%	14.30	9.87	12.04
p=75%	13.54	9.20	11.15
p=100%	13.81	9.45	11.24

Table 3: Ablation Study of the SS-MPC based on masking rate during post-training.

Model	Target Utt.	Adr.
BERT (Devlin et al., 2019)	-	82.88%
SA-BERT (Sun et al., 2019)	-	86.98%
MPC-BERT (Gu et al., 2021)	-	89.54%
GIFT (Gu et al., 2023a)	-	90.18%
SS-MPC_{Encoder}	76.38%	89.92%

Table 4: Performance (precision@1) of predicting target-utterance and addressee on the test set of Ubuntu IRC (2019) (Hu et al., 2019).

domly sampled 100 conversations from the Ubuntu IRC benchmark dataset and asked three graduate students to evaluate the quality of the generated responses. The evaluation focuses on three independent aspects: (1) relevance, (2) fluency, and (3) informativeness. Each judge assigned binary scores for each aspect, with the final score ranging from 0 to 3. Table 2 presents the average final score of human evaluation results comparing GPT-2, BART, MADNet, and SS-MPC against the ground truth (Gold Label). In addition, Fleiss’s Kappa (Fleiss, 1971) is calculated to measure inter-annotator agreement. The result indicates that SS-MPC produces more relevant, fluent, and informative responses than any other model.

4.3 Ablation Study

Post-Training and Masking Probability Table 3 shows the performance of the model as a function of the probability p of masking the target-utterance

Models	Structural Info.	BLEU-1	BLEU-2	BLEU-3	BLEU-4	METEOR	ROUGE-L
MADNet (Gu et al., 2023b)	fully given	11.69	4.57	2.33	1.45	9.48	11.82
SS-MPC		15.60	6.62	3.67	2.44	10.92	12.44
SS-MPC _{real-world}	-last utt.	13.59	5.47	3.10	2.14	9.15	11.03

Table 5: Performance of SS-MPC without structural information of last utterance.

Model	Utt. Info.	Target Utt.	Adr.
GIFT (Gu et al., 2023a)	-	-	90.18%
SS-MPC _{Encoder}		76.38%	89.92%
SS-MPC _{real-world}	-last utt.	62.90%	78.40%

Table 6: Performance of predicting target-utterance of addressee in other MPC model. Unlike other models, SS-MPC Encoder does not utilize final response.

and addressee information during post-training of the encoder. From Table 3, we can see that the post-trained model with 25% of masking probability effects impressively on final performance of generation, while 50% of masking probability achieves almost same performance as the model only fine-tuned without post-training. The post-trained models with 75% and 100% of masking probability shows even lower performances than the model only fine-tuned without post-training, which can be interpreted as excessive masking degrades the model’s analytical capability. This results provides the partial criteria of required masking probability when re-learning newly added token embeddings using the Masked Language Modeling (MLM) approach.

Addressee prediction The SS-MPC is trained to understand structures through structure tokens via post-training process. We hypothesize that this approach could be applied to the task of addressee prediction. To verify this, we train the model by masking only the addressee and predicting it. Table 4 shows the comparison of the addressee prediction tasks with other MPC analysis models, in terms of Precision@1. It shows that SS-MPC achieves 89.92% in addressee prediction, which is very close to the previous SOTA model, GIFT.

4.4 Response Generation in Real-World MPC Scenario

SS-MPC can generate responses even when partial structural information is missing by simply masking the absent tokens as Equation 6. Here, we assume the real-world MPC scenario, where the model should continuously generate the MPC. In this scenario, the target-utterance and addressee information of the response is missing. And the

model has to predict the missing target-utterance and addressee while generating the following response. Our method can accumulate the predicted structure information to generate the next response continuously in this scenario. "[IDX_t]" and "[SPK_a]" are masked with structure masking tokens to construct sequence-structured input (Equation 8).

Table 5 presents the performance of SS-MPC without target-utterance and addressee information for the final response, which is marked as SS-MPC_{real-world} in this table. While its performance is slightly lower than SS-MPC with full structural information, it still outperforms the previous SOTA model, MADnet, in BLEU and maintains comparable scores in METEOR and ROUGE-L. Unlike existing SOTA models that require all the target-utterance and addressee information for each utterance, SS-MPC can generate the response with predicting the most appropriate target-utterance and addressee to answer itself.

In addition, Table 6 demonstrates SS-MPC’s ability to predict the target-utterance and addressee of the last utterance in the real-world scenario. Although the performance is inevitably lower than in an environment where responses are present, SS-MPC still appears to maintain a reasonable performance on the target-utterance and addressee prediction in this scenario.

5 Conclusions

We introduce SS-MPC, a model optimized for generating responses in multi-party conversations. Unlike traditional graph-based approaches, SS-MPC employs the encoder-decoder architecture of transformer to fully leverage the pre-trained knowledge of language models. For this, we propose a novel method to encode dialogue structure sequentially within the input, allowing the model to capture the interaction flow in the dialogue without relying on explicit graph representations. SS-MPC outperforms the existing SOTA MPC response generation model and has the distinct advantage of easily application to real-world MPC scenarios depending on not requiring any additional module.

Limitations

The SS-MPC proposed in this paper has shown good performance compared to existing MPC response generation models, but it is limited by the fact that both training and inference are performed only on the Ubuntu IRC benchmark dataset, which makes it less generalizable. This is due to the absolute lack of MPC datasets, and it is necessary to apply the model to a wider variety of topics and conversations between different speakers to maintain generality. There is another room for further development of the model. For example, the model can be trained by initializing the initial embeddings of the added soft prompt tokens to specific values (e.g., [CLS] or [SEP] embeddings), or by initializing the embeddings to follow a specific distribution. Acquiring additional MPC datasets and further developing Multi-MPC will be part of our future work.

References

- Omer Anjum, Chak Ho Chan, Tanitpong Lawphongpanich, Yucheng Liang, Tianyi Tang, Shuchen Zhang, Wen mei W. Hwu, Jinjun Xiong, and Sanjay J. Patel. 2020. [Vertex: An end-to-end ai powered conversation management system for multi-party chat platforms](#). *Companion Publication of the 2020 Conference on Computer Supported Cooperative Work and Social Computing*.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. [BERT: Pre-training of deep bidirectional transformers for language understanding](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186, Minneapolis, Minnesota. Association for Computational Linguistics.
- Joseph Fleiss. 1971. [Measuring nominal scale agreement among many raters](#). *Psychological Bulletin*, 76:378–.
- Jia-Chen Gu, Zhen-Hua Ling, QUAN LIU, Cong Liu, and Guoping Hu. 2023a. [Gift: Graph-induced fine-tuning for multi-party conversation understanding](#). In *Annual Meeting of the Association for Computational Linguistics*.
- Jia-Chen Gu, Chao-Hong Tan, Caiyuan Chu, Zhen-Hua Ling, Chongyang Tao, Quan Liu, and Cong Liu. 2023b. [MADNet: Maximizing addressee deduction expectation for multi-party conversation generation](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 7681–7692, Singapore. Association for Computational Linguistics.
- Jia-Chen Gu, Chao-Hong Tan, Chongyang Tao, Zhen-Hua Ling, Huang Hu, Xiubo Geng, and Daxin Jiang. 2022. [HeterMPC: A heterogeneous graph neural network for response generation in multi-party conversations](#). In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 5086–5097, Dublin, Ireland. Association for Computational Linguistics.
- Jia-Chen Gu, Chongyang Tao, Zhenhua Ling, Can Xu, Xiubo Geng, and Daxin Jiang. 2021. [Mpc-bert: A pre-trained language model for multi-party conversation understanding](#). *ArXiv*, abs/2106.01541.
- Wenpeng Hu, Zhangming Chan, Bing Liu, Dongyan Zhao, Jinwen Ma, and Rui Yan. 2019. [Gsn: A graph-structured network for multi-party dialogues](#). In *International Joint Conference on Artificial Intelligence*.
- Ran Le, Wenpeng Hu, Mingyue Shang, Zhenjun You, Lidong Bing, Dongyan Zhao, and Rui Yan. 2019. [Who is speaking to whom? learning to identify utterance addressee in multi-party conversations](#). 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 1909–1919, Hong Kong, China. Association for Computational Linguistics.
- Mike Lewis, Yinhan Liu, Naman Goyal, Marjan Ghazvininejad, Abdelrahman Mohamed, Omer Levy, Veselin Stoyanov, and Luke Zettlemoyer. 2020. [BART: Denoising sequence-to-sequence pre-training for natural language generation, translation, and comprehension](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 7871–7880, Online. Association for Computational Linguistics.
- Yiyang Li and Hai Zhao. 2023. [EM pre-training for multi-party dialogue response generation](#). In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 92–103, Toronto, Canada. Association for Computational Linguistics.
- Ilya Loshchilov and Frank Hutter. 2017. [Decoupled weight decay regularization](#). In *International Conference on Learning Representations*.
- Ryan Lowe, Nissan Pow, Iulian Serban, and Joelle Pineau. 2015. [The Ubuntu dialogue corpus: A large dataset for research in unstructured multi-turn dialogue systems](#). In *Proceedings of the 16th Annual Meeting of the Special Interest Group on Discourse and Dialogue*, pages 285–294, Prague, Czech Republic. Association for Computational Linguistics.
- Shikib Mehri, Evgeniia Razumovskaia, Tiancheng Zhao, and Maxine Eskenazi. 2019. [Pretraining methods for dialog context representation learning](#). In *Proceedings of the 57th Annual Meeting of the Association for*

Computational Linguistics, pages 3836–3845, Florence, Italy. Association for Computational Linguistics.

Zhao Meng, Lili Mou, and Zhi Jin. 2017. [Towards neural speaker modeling in multi-party conversation: The task, dataset, and models](#). In *International Conference on Language Resources and Evaluation*.

Hiroki Ouchi and Yuta Tsuboi. 2016. [Addressee and response selection for multi-party conversation](#). In *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing*, pages 2133–2143, Austin, Texas. Association for Computational Linguistics.

Sunjeong Park and Youn-kyung Lim. 2020. [Investigating user expectations on the roles of family-shared ai speakers](#). In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, CHI ’20, page 1–13, New York, NY, USA. Association for Computing Machinery.

Alec Radford, Jeff Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. 2019. Language models are unsupervised multitask learners.

Weizhou Shen, Junqing Chen, Xiaojun Quan, and Zhixiang Xie. 2020. [Dialogxl: All-in-one xlnet for multi-party conversation emotion recognition](#). In *AAAI Conference on Artificial Intelligence*.

Chi Sun, Luyao Huang, and Xipeng Qiu. 2019. [Utilizing BERT for aspect-based sentiment analysis via constructing auxiliary sentence](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 380–385, Minneapolis, Minnesota. Association for Computational Linguistics.

Weishi Wang, Steven C.H. Hoi, and Shafiq Joty. 2020. [Response selection for multi-party conversations with dynamic topic tracking](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 6581–6591, Online. Association for Computational Linguistics.

Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Remi Louf, Morgan Funtowicz, Joe Davison, Sam Shleifer, Patrick von Platen, Clara Ma, Yacine Jernite, Julien Plu, Canwen Xu, Teven Le Scao, Sylvain Gugger, Mariama Drame, Quentin Lhoest, and Alexander Rush. 2020. [Transformers: State-of-the-art natural language processing](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, pages 38–45, Online. Association for Computational Linguistics.

Rui Zhang, Honglak Lee, Lazaros Polymenakos, and Dragomir Radev. 2018. Addressee and response selection in multi-party conversations with speaker interaction rnns. In *Proceedings of the Thirty-Second AAAI Conference on Artificial Intelligence*

and Thirtieth Innovative Applications of Artificial Intelligence Conference and Eighth AAAI Symposium on Educational Advances in Artificial Intelligence, AAAI’18/IAAI’18/EAAI’18. AAAI Press.

A Data Statistics

Datasets	Train	Valid	Test
Ouchi and Tsuboi (2016)	461,120	28,570	32,668
Hu et al. (2019)	311,725	5,000	5,000

Table 7: Statistics of the two benchmarks evaluated in this paper.

B Case Study

Case 1

SPK1: e : could not open lock file FILEPATH - open (13 permission denied) e : unable to lock the administration directory (FILEPATH are you root ?

SPK2: you need to be using sudo

SPK3: ah , ok , so it 's not just me

SPK4: you can read all about FILEPATH issues on URL

SPK2: you can edit off the us . part of the url and that should fix you up

target-utterance: e : could not open lock file FILEPATH - open (13 permission denied) e : unable to lock the administration directory (FILEPATH are you root ?

answer speaker: SPK5

answer addressee: SPK1

Generated Responses:

Models	Response
Gold Label	well , you are n't root , are n't you ? EMOJI
Bart	are you using sudo or gksudo ?
HeterMPC	please don't flood, use url to paste, do nix, do not use enter as punctuation.
MADNet	please don't flood, use URL to paste, do nt use enter as punctuation.
SS-MPC	you need to be able to open it as root

Case 2

SPK1: many iconpacks may be for the older kde3

SPK2: i have run checkdisk from windows and i have many errors in partition

SPK3: sure , ai n't there a way to filter out 4.1 packs ?

SPK1: thats not a good sign .

SPK3: i 'm using kde to do that

target-utterance: sure , ai n't there a way to filter out 4.1 packs ?

answer speaker: SPK1

answer addressee: SPK3

Generated Responses:

Models	Response
Gold Label	i imagine it depends on FILEPATH you are looking
Bart	i 'm not sure , but i guess it 's filter by version
HeterMPC	i don't know, i'm not sure how to do it
MADNet	i don't know, i've never used kde.
SS-MPC	there is a way , but i do n't remember the kde3 way

Case 3

SPK1: shaft : not now , only with the ubuntu cd i can see the disk unpartitioned
SPK2: see when i stopped a linux distro during install a while back everything refused to boot from cd until i sorted it out
SPK1: i guess is something like that . how do i do that ?
SPK2: well its just a guess from previous experience ,
SPK2: id try installing ubuntu , with it using all your hdd , then try installing windows over it , thats what i had to do
SPK1: you mean , i will format the disk again with windows , errasing the previous ubuntu installing ?

target-utterance: you mean , i will format the disk again with windows , errasing the previous ubuntu installing ?
answer speaker: SPK2
answer addressee: SPK1

Generated Responses:

Models	Response
Gold Label	going off my experience with what seemed the same problem yes , unless you can ask someone else that may have a technical solution for you EMOJI
Bart	i will format the disk again with windows , and erase the previous ubuntu install
HeterMPC	i think you can do that with the ubuntu cd
MADNet	i dont know, i dont use windows, so i dont know
SS-MPC	you should be able to use the entire drive again , but i dont know if it will work

C Adaptation on Large Language Models

Models	BLEU-1	METEOR	ROUGE-L
Qwen2.5-3B	33.62	31.89	33.09
SS-MPC_{Qwen2.5-3B}	34.62	34.20	33.33

Table 8: Applied result on LLM. We use Qwen2.5-3B model for training and inference.

The concept of dialogue structuralization is also applicable to Large Language Model (LLM). In Table 8, we compare the effects of dialogue structuralization in LLMs. We adopt sequence-structured input in Qwen2.5-3B. The results demonstrate that sequence-structured input significantly impacts performance. Especially, Qwen achieved an improvement of nearly 1%p in BLEU-4 and 2.5%p in METEOR solely by utilizing sequence-structured input. This highlights the importance of incorporating conversational structure into the input representation.

D License

The data used in this paper can be found on <https://github.com/ryanzhumich/sirnn> (Ubuntu IRC (2016)) and <https://github.com/morning-dews/GSN-Dialogues> (Ubuntu IRC (2019)) We utilize parts of the code provided by HeterMPC⁵, which is licensed under the Apache 2.0 License.

⁵<https://github.com/lxchtan/HeterMPC>