

HyCoRA: Hyper-Contrastive Role-Adaptive Learning for Role-Playing

Shihao Yang^{1*}, Zhicong Lu^{2*}, Yong Yang^{1*}, Bo Lv², Yang Shen¹, Nayu Liu^{1†}

¹Tianjin Laboratory Autonomous Intelligence Technology and Systems,
School of Computer Science and Technology, Tiangong University

²University of Chinese Academy of Sciences
{yshihao, nayuliu}@tiangong.edu.cn

Abstract

Multi-character role-playing aims to equip models with the capability to simulate diverse roles. Existing methods either use one shared parameterized module across all roles or assign a separate parameterized module to each role. However, the role-shared module may ignore distinct traits of each role, weakening personality learning, while the role-specific module may overlook shared traits across multiple roles, hindering commonality modeling. In this paper, we propose a novel **HyCoRA: Hyper-Contrastive Role-Adaptive** learning framework, which efficiently improves multi-character role-playing ability by balancing the learning of distinct and shared traits. Specifically, we propose a Hyper-Half Low-Rank Adaptation structure, where one half is a role-specific module generated by a lightweight hyper-network, and the other half is a trainable role-shared module. The role-specific module is devised to represent distinct persona signatures, while the role-shared module serves to capture common traits. Moreover, to better reflect distinct personalities across different roles, we design a hyper-contrastive learning mechanism to help the hyper-network distinguish their unique characteristics. Extensive experimental results on both English and Chinese available benchmarks demonstrate the superiority of our framework. Further GPT-4 evaluations and visual analyses also verify the capability of HyCoRA to capture role characteristics.

Code — <https://github.com/yshihao-ai/HyCoRA>

Introduction

The roles in books and movies allow us to witness rich personalities and narratives (Lu et al. 2023). However, we are unable to deeply interact with them in real time, which limits opportunities for immersive experiences with roles. Therefore, developing an AI system that can effectively imitate various roles is meaningful. To meet this need, multi-character role-playing (MCRP) (Wu et al. 2025; Tang et al. 2025; Castricato et al. 2025; Tseng et al. 2024) has been proposed to deliver an immersive conversational experience in real time by enabling models to simulate diverse characters.

*These authors contributed equally.

†Corresponding author.

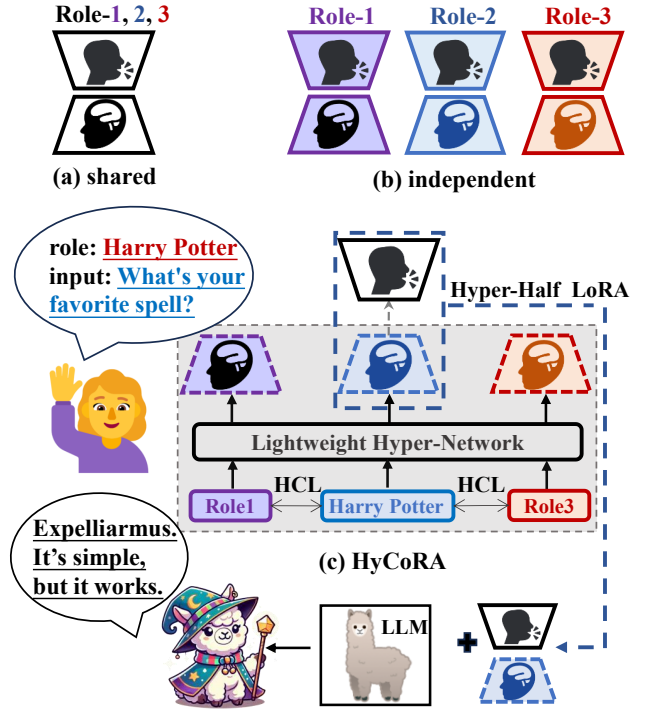


Figure 1: (a) and (b) denote strategies employing a shared module across roles and independent modules assigned to each role, respectively. (c) refers to the HyCoRA proposed in this paper. HCL: Hyper-Contrastive Learning.

The instruction-following capability of large language models (LLMs) (Grattafiori et al. 2024; Yang et al. 2024; Jiang et al. 2023; Wei et al. 2025) has enabled them to show promising performance in MCRP. Considering the limited amount of role-playing data, recent approaches often adopt the low-rank adaptation (LoRA) (Hu et al. 2021) method to enhance role-playing ability while preserving the general capabilities of the pre-trained model. As shown in Figure 1, depending on how they capture role characteristics, existing methods can be broadly grouped into two paradigms: (1) Using a shared module for all characters (Tao et al. 2024; Wang et al. 2024a); and (2) Assigning an independent mod-

ule to each character (Shao et al. 2023; Yu et al. 2024). Both paradigms have shown clear advantages in capturing role behaviors, but some challenges still remain and warrant further investigation. Role-shared module may limit the model’s ability to capture fine-grained role-specific traits, due to interference among diverse character features within a single representation space (Baziotis et al. 2022). Role-specific module not only limits the transfer of common features across roles but may also suffer from data sparsity for each adapter, leading to insufficient training, while additionally incurring significant parameter overhead.

In this paper, we propose a HyCoRA: Hyper-Contrastive Role-Adaptive Learning framework, which enhances multi-character role-playing ability by balancing the learning of distinct and shared characteristics. Intuitively, different roles interpret questions from distinct views and give answers following shared linguistic rules and common knowledge. To capture both unique and shared behaviors, we introduce a Hyper-Half LoRA structure, as shown in Figure 1c. In this structure, we decouple the low-rank adaptation matrices A and B in the LoRA module to serve different functions. Matrix A is generated by a lightweight hyper-network to capture role-specific traits, while Matrix B is a shared trainable matrix to encode common features. Inspired by the hyper-network technique, instead of assigning an independent matrix A to each role, we generate it through a lightweight hyper-network to improve generalization with limited data per role and reduce overall parameters as the number of roles grows. Additionally, to better reflect distinct personalities across different roles, we design a hyper-contrastive learning mechanism to help the hyper-network distinguish their unique traits. It encourages the model to effectively learn each role’s traits from its own responses while distinguishing them from traits hidden in other roles’ responses.

Extensive experimental results on both English and Chinese available benchmarks demonstrate the superiority of our framework. Furthermore, GPT-4 evaluations and visual analyses indicate HyCoRA’s ability to capture distinct characteristics. In summary, our main contributions include:

- We propose a HyCoRA: Hyper-Contrastive Role-Adaptive Learning framework, which efficiently improves multi-character role-playing ability by balancing the learning of distinct and shared characteristics.
- HyCoRA adopts a Hyper-Half LoRA structure that incorporates a role-specific persona projection module generated by a lightweight hyper-network and a role-shared trait representation module. In addition, we design a hyper-contrastive learning mechanism to better distinguish role-specific characteristics.
- Extensive experiments show that HyCoRA achieves strong performance across both Chinese and English scenarios, as validated by automatic and model-based evaluations. Additionally, we conduct a visualization analysis to provide further insights into the Hyper-Half LoRA structure. The related code will be released to facilitate future research.

Related Work

Multi-Character Role-Playing

Multi-character role-playing (MCRP) (Chen et al. 2024a; Tu et al. 2024; Ran et al. 2024; Chen et al. 2023) enables models to emulate diverse roles, allowing them to display various behaviors. Recent research in MCRP has made significant progress (Wang et al. 2024b; Wu et al. 2024). Some methods share parameters across multiple roles, such as DITTO (Lu et al. 2024), Character-GLM (Zhou et al. 2024), RoleCraft-GLM (Tao et al. 2024), and RoleLLM (Wang et al. 2024a). These methods can effectively capture common traits shared by different roles, improving generalization. Other methods allocate distinct parameters to each role, such as Character-LLM (Shao et al. 2023) and Neeko (Yu et al. 2024), which helps maintain each role’s unique personality. To balance model size and learning ability, both types of methods often use Low-Rank Adaptation (LoRA) techniques (Hu et al. 2021) to effectively simulate diverse roles. Despite the improvements made by these methods, there remain several challenges to overcome. Methods that share parameters may overlook the unique traits of each role, which may weaken personality learning. In contrast, methods with distinct parameters may ignore common traits across roles, limiting the modeling of shared knowledge. In this work, we address these challenges by balancing shared and distinct traits learning to enhance MCRP capabilities.

Hyper-Network

Hyper-networks (Liu et al. 2025a; Lv et al. 2024; Li et al. 2024; Baziotis et al. 2022) have been widely adopted in multi-task (Iverson et al. 2023; Liu et al. 2024, 2025c; Wei et al. 2023; Liu et al. 2025b) learning scenarios. Motivated by their ability to dynamically generate parameters for target networks, we view MCRP as a multi-task learning problem, with each role treated as a distinct task. In HyCoRA, rather than training independent parameters for each role with limited role-specific data, the hyper-network is trained across data from all roles, which not only reduces parameter overhead but also addresses insufficient training caused by limited data for each role. Moreover, traditional hyper-networks incur prohibitive computational costs when applied to LLMs. To address this issue, we propose a lightweight hyper-network that notably reduces parameter cost. Additionally, we employ hyper-contrastive learning to help the lightweight hyper-network distinguish features of different roles.

Methodology

In this section, we detail our proposed HyCoRA framework for MCRP, which balances the learning of distinct and shared traits. As shown in Figure 2a, HyCoRA adopts a Hyper-Half LoRA structure, which consists of a role-specific persona projection module generated by a lightweight hyper-network and a role-shared trait representation module implemented as a trainable matrix. The role-specific module aims to represent distinct persona signatures, while the role-shared module serves to capture common features. Incorporating this structure into the backbone

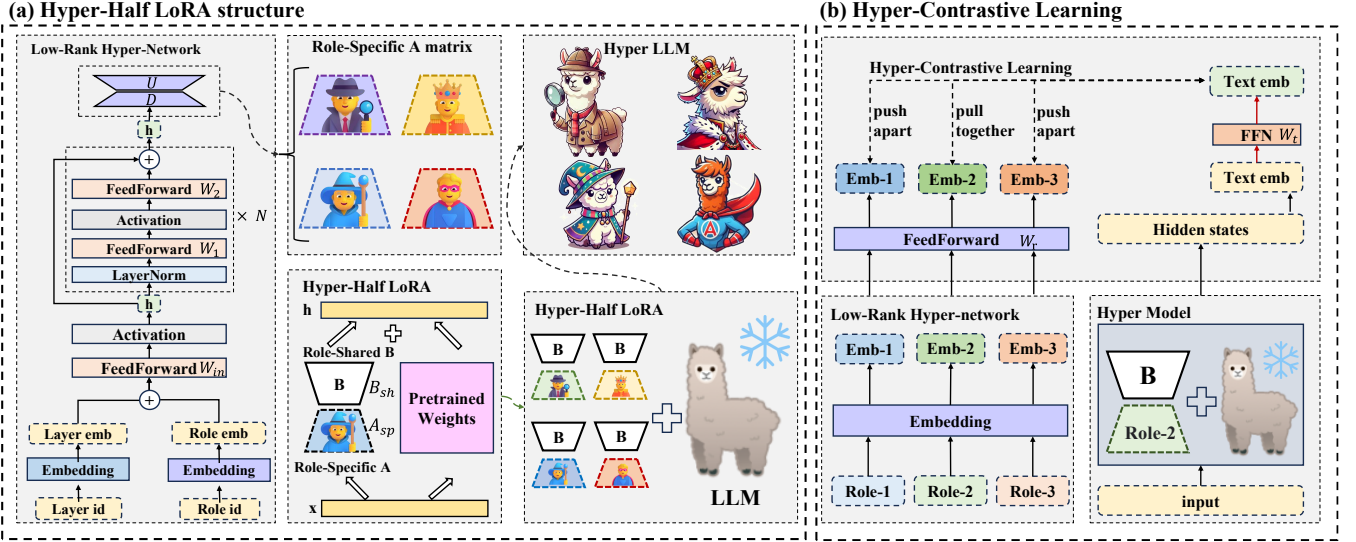


Figure 2: An illustration of our framework for MCRP. (a) We construct the Hyper-Half LoRA structure, where the role-specific matrix A is generated by a lightweight hyper-network, and the role-shared matrix B is implemented as a trainable matrix. (b) We introduce a hyper-contrastive learning mechanism that pulls role representations closer to response representations from the same role and pushes them away from those of different roles.

LLM yields the Hyper Model. During training, as shown in Figure 2b, we employ hyper-contrastive learning as an auxiliary task, which utilizes role-specific responses to further help the lightweight hyper-network capture and distinguish distinct characteristics.

Hyper-Half LoRA

Role-Specific A The lightweight hyper-network leverages character IDs to generate role-specific matrices A. Additionally, since the matrices should differ across layers even for the same character, layer IDs are introduced to generate layer-specific matrices, as shown in Figure 2a. We use c and l to denote the character ID and layer ID, respectively. The character and layer embeddings, $e_c \in \mathbb{R}^{d_c}$ and $e_l \in \mathbb{R}^{d_l}$, are obtained by embedding c and l using separate embedding matrices. Then, we concatenate both embeddings and project them with $W_{in} \in \mathbb{R}^{d_h \times (d_c + d_l)}$ to obtain the hyper-network hidden state $h_0 \in \mathbb{R}^{d_h}$. Finally, we pass h_0 through N residual blocks to get $h_i \in \mathbb{R}^{d_h}$, which aims to capture high-level features:

$$h_0 = \text{ReLU}(W_{in}([e_c; e_l]) + b_{in}), \quad (1)$$

$$h_{i+1} = W_2(\text{ReLU}(W_1 h_i + b_1)) + b_2, \quad (2)$$

where $W_1 \in \mathbb{R}^{d_h \times d_h}$ and $W_2 \in \mathbb{R}^{d_h \times d_h}$ are the trainable weights of each residual block.

In traditional hyper-networks, the final hidden state h_N is projected by $W_A \in \mathbb{R}^{(r \times d) \times d_h}$, where (r, d) corresponds to the shape of the role-specific matrix A_{sp} . However, as d_h increases, the parameter size of W_A grows significantly. To address this challenge, we decompose W_A into two low-rank matrices: $D \in \mathbb{R}^{k \times d_h}$ and $U \in \mathbb{R}^{(r \times d) \times k}$, where $k \ll \min(d_h, r \times d)$. The role-specific matrix A_{sp} is computed

as follows:

$$A_{sp} = U(Dh_N + b_d) + b_u. \quad (3)$$

As d_h increases, the size of U remains unchanged, effectively limiting the overall parameter growth.

The generation process of role-specific matrices A illustrates that the lightweight hyper-network is trained on data from all characters, which ensures sufficient training of the hyper-network. Additionally, in order to generate informative and discriminative role-specific matrices A, the role embedding e_c is required to model the unique traits and distinguish between different characters. Therefore, we will use the hyper-contrastive learning mechanism at a later stage to help the lightweight hyper-network better differentiate between the roles.

Role-Shared B In contrast to the Role-Specific matrix $A_{sp} \in \mathbb{R}^{r \times d}$, the Role-Shared matrix $B_{sh} \in \mathbb{R}^{m \times r}$ is a trainable matrix, rather than being generated by the lightweight hyper-network. Similar to LoRA, we initialize separate trainable matrices B_{sh} for different layers. However, within the same layer, multiple roles share a single matrix B_{sh} . Let $W_0 \in \mathbb{R}^{m \times d}$ represent the weight matrix of a pretrained LLM. For the original $out = W_0 x$, the modified forward pass in the Hyper-Half LoRA structure is given by:

$$out = W_0 x + \frac{\alpha}{r} B_{sh} A_{sp} x, \quad (4)$$

where α is a scaling factor, and r is the decomposition rank.

Multiple roles share the matrix B_{sh} , which not only reduces the number of trainable parameters compared to fully role-specific decompositions, but also helps the model capture common features across roles. When integrated into the Hyper-Half LoRA structure, this role-shared component allows the lightweight hyper-network to focus on modeling role-specific traits via A_{sp} .

Hyper-Contrastive Learning

A role’s preferences and personality are often implicitly reflected in its responses. Since each role is represented by an embedding in the lightweight hyper-network, these responses can serve as useful signals to help the hyper-network learn more discriminative role embeddings. Motivated by this, we propose a hyper-contrastive learning mechanism, which leverages character responses to help the lightweight hyper-network better distinguish role-specific traits.

We refer to an LLM equipped with the Hyper-Half LoRA structure as a Hyper Model. As illustrated in Figure 2b, for each batch, role IDs are embedded via the lightweight hyper-network to produce role embeddings. Simultaneously, the Hyper Model processes the question-answer pairs to generate token-level hidden states. To capture the overall semantics of the response, we extract the hidden state of the final token as the sequence representation. This final state and the role embeddings are separately projected into an m -dimensional space, yielding the role and response representations:

$$z_{r_i} = W_r e_{c_i} + b_r, \quad z_{t_i} = W_t e_{t_i} + b_t, \quad (5)$$

where e_{c_i} is the role embedding of the role in the i -th sample, e_{t_i} is the final token’s hidden state of the corresponding response, and W_r, W_t, b_r, b_t are trainable projection parameters.

To encourage the lightweight hyper-network better distinguish role-specific characteristics, we define a hyper-contrastive learning loss. For the i -th sample, the role representation is encouraged to be close to the response representations in the batch that have the same role as i -th sample (i.e., positive samples), and to be distant from those of samples with different roles (i.e., negative samples). The cross-entropy loss and the hyper-contrastive loss together constitute the HyCoRA’s loss function. However, due to the randomized state of the hyper-network parameters in the early stages of training, the final token hidden state cannot effectively represent the text sequence. Therefore, only the cross-entropy loss is initially used. Once the model begins to capture basic sentence semantics, the hyper-contrastive loss is introduced. The overall loss function is defined as:

$$\mathcal{L} = - \sum_{i=1}^n \log P(y_i | x_{\leq i}) + \lambda \mathcal{L}_{cl}, \quad (6)$$

$$\mathcal{L}_{cl} = \sum_{i \in I} \frac{-1}{|P(i)|} \sum_{p \in P(i)} \log \frac{e^{\text{sim}(z_{r_i}, z_{t_p})/\tau}}{\sum_{j \in I} e^{\text{sim}(z_{r_i}, z_{t_j})/\tau}}, \quad (7)$$

$$\lambda = \lambda_{\min} + (\lambda_{\max} - \lambda_{\min}) \cdot \max \left(0, \frac{ep - ep_{st}}{ep_{to} - ep_{st}} \right), \quad (8)$$

where $\text{sim}(p, q)$ denotes cosine similarity, I is the set of samples in a batch, and $P(i)$ represents the subset of samples in I that have the same role as the i -th sample. z_{r_i} and z_{t_p} are the role and response representations, respectively. λ controls the weight of the hyper-contrastive loss, increasing linearly from $\lambda_{\min}$ to $\lambda_{\max}$ between epoch ep_{st} and ep_{to} .

Experiments

Datasets

To verify the effectiveness of our framework in MCRP, we use RoleBench (Wang et al. 2024a) as the benchmark dataset. RoleBench contains 168,093 samples across 100 distinct roles (5 in Chinese and 95 in English), covering both general and role-specific knowledge. General knowledge assesses the model’s ability to capture tone and personality, while role-specific knowledge evaluates its understanding of the character’s background.

Implementation Details

For comprehensive evaluation across backbone models, we adopt several popular LLMs. For the Chinese scenario, we use ChatGLM2-6B¹ (Zeng et al. 2022), Qwen2-7B² (Bai et al. 2023), and Qwen2.5-7B³ (Yang et al. 2024) as backbone models. For the English scenario, we adopt LLaMA-2-7B⁴ (Touvron et al. 2023) as the backbone model. Specifically, LLaMA-2-7B, Qwen2-7B, and Qwen2.5-7B are pre-trained models, while ChatGLM2-6B is post-trained with enhanced dialogue capabilities.

Evaluation metrics

To evaluate our framework from multiple perspectives, we employ both automatic and model-based evaluation methods. Following previous research (Chen et al. 2024b; Tao et al. 2024; Wang et al. 2024a; Yu et al. 2024), we use ROUGE (Lin 2004) and BLEU (Papineni et al. 2002) as automatic metrics to quantitatively assess output quality. For model-based evaluation, we use a ranking-based approach to calculate win rates. Specifically, GPT-4 assesses the model’s performance across two dimensions: **(1) Role Behavior (RB)**, which examines whether the responses align with the character’s tone and personality traits, and **(2) Role Knowledge (RK)**, which evaluates the consistency of responses with the character’s background knowledge.

Baselines

(1) RoleLLM (Wang et al. 2024a) adopts a shared LoRA module across roles, and the same strategy is also used in RoleCraft-GLM (Tao et al. 2024). (2) Independent LoRA (Ind. LoRA) assigns an independent LoRA module to each role. (3) Neeko (Yu et al. 2024) adopts a similar approach to Ind.LoRA but introduces an additional gating mechanism. Due to cost constraints, (2) and (3) are only compared in the Chinese scenario. (4) Role-playing models are used as baselines due to their strong role-playing capabilities, including ChatPLUG (Tian et al. 2023) and Character.AI⁵. (5) Instruction-tuned models are used as baselines due to their excellent instruction-following capabilities. These models include LLaMA-2-7B-Chat (Touvron et al. 2023), Vicuna-13B (Chiang et al. 2023), and Alpaca-7B (Taori et al. 2023)

¹<https://huggingface.co/THUDM/chatglm2-6b>

²<https://huggingface.co/Qwen/Qwen2-7B>

³<https://huggingface.co/Qwen/Qwen2.5-7B>

⁴<https://huggingface.co/meta-llama/Llama-2-7b-hf>

⁵<https://character.ai/>

Models	General				Specific			
	BLEU ↑	ROUGE-1 ↑	ROUGE-2 ↑	ROUGE-L ↑	BLEU ↑	ROUGE-1 ↑	ROUGE-2 ↑	ROUGE-L ↑
GPT-4	37.02	58.55	36.53	52.14	07.32	35.57	12.28	23.06
ChatPLUG	16.14	40.43	19.50	34.63	08.14	36.33	13.44	24.53
Character.AI	15.76	44.92	20.71	36.45	08.74	37.17	13.91	25.58
Yi-6B-Chat	16.55	40.21	19.82	34.17	06.74	33.03	11.68	21.99
ChatGLM2	14.83	39.48	18.73	32.91	07.74	36.08	12.82	23.15
ChatGLM2-LoRA (RoleLLM)	34.45	56.48	34.29	50.22	12.09	42.62	18.66	30.14
ChatGLM2-HyCoRA	35.09	56.65	34.56	50.69	13.40	43.11	19.78	30.71
Qwen2-Instruct	17.25	43.98	20.30	36.14	05.20	32.13	08.95	20.21
Qwen2-LoRA (RoleLLM)	35.54	56.85	34.84	50.60	13.86	44.19	20.63	31.92
Qwen2-HyCoRA	35.80	57.16	35.38	51.18	13.88	44.34	20.50	32.02
Qwen2.5-Instruct	19.07	46.07	22.27	37.68	04.63	31.25	08.37	19.43
Qwen2.5-LoRA (RoleLLM)	34.18	56.21	34.01	50.25	14.15	44.07	20.82	31.65
Qwen2.5-HyCoRA	34.80	56.78	34.49	50.65	14.21	44.23	20.92	32.51

Table 1: Model performance assessment in the Chinese scenario based on automatic evaluation.

Models	General				Specific			
	BLEU ↑	ROUGE-1 ↑	ROUGE-2 ↑	ROUGE-L ↑	BLEU ↑	ROUGE-1 ↑	ROUGE-2 ↑	ROUGE-L ↑
GPT-4	42.67	55.03	32.62	53.56	08.01	32.04	09.89	28.83
ChatPLUG	10.00	24.01	10.57	22.57	04.24	25.45	05.48	22.67
Character.AI	17.66	30.03	14.46	28.18	05.05	25.74	06.45	23.27
Vicuna-13B	18.67	34.88	15.80	33.50	06.30	29.05	08.14	26.06
Alpaca-7B	13.52	27.58	13.66	26.79	07.06	31.17	09.87	27.89
LLaMA-2-7B-Chat	11.72	28.60	12.21	26.88	02.33	20.55	04.29	18.81
LLaMA-2-LoRA (RoleLLM)	25.50	38.68	19.97	37.11	09.95	34.63	12.02	31.14
LLaMA-2-HyCoRA	26.26	39.29	21.23	37.65	10.12	34.85	12.15	31.31

Table 2: Model performance assessment in the English scenario based on automatic evaluation.

for the English scenario, as well as ChatGLM2-6B, Yi-6B-Chat⁶, Qwen2-7B-Instruct (Bai et al. 2023), and Qwen2.5-7B-Instruct (Yang et al. 2024) for the Chinese scenario. (6) GPT-4 is also selected to serve as an upper-bound reference model.

Performance Analysis

Automatic Evaluation Results are reported in Tables 1 and 2. Overall, HyCoRA outperforms other baselines, reflecting its superior ability to simulate diverse roles. Compared with role-playing and instruction-tuned models, LoRA and HyCoRA demonstrate notable advantages, which reveal the challenges of capturing diverse character traits in multi-character role-playing. Compared with LoRA, HyCoRA shows more marked advantages in the Chinese scenario than in the English scenario. We attribute this to HyCoRA generating role-specific matrices directly using role IDs, without requiring role-specific personality instructions, which lowers token consumption. This indicates that HyCoRA learns character traits from role responses. However, as the number of roles increases, similar responses across roles become inevitable, which may confuse the model. Therefore, we spec-

ulate that providing HyCoRA with role-specific personality instructions, similar to those used in LoRA, could potentially enhance its performance. Furthermore, HyCoRA outperforms GPT-4 in some scenarios but falls short in others, which may be attributed to GPT-4’s large parameter size.

Methods	Datasets	Role-1	Role-2	Role-3	Role-4	Role-5
Ind. LoRA	General	50.30	42.85	49.46	49.81	48.38
Neeko	General	51.02	42.31	52.35	50.32	48.32
HyCoRA	General	53.58	43.73	54.33	52.66	50.13
Ind. LoRA	Specific	25.46	27.04	30.40	27.77	31.97
Neeko	Specific	30.41	28.22	31.20	27.72	31.12
HyCoRA	Specific	32.56	31.53	34.19	29.51	30.88

Table 3: ROUGE-L evaluation results for Independent LoRA, Neeko, and HyCoRA.

Comparison with Ind. LoRA and Neeko. To further assess model performance, we conduct experiments to evaluate the imitation ability of each role, as shown in Table 3. Specifically, we report ROUGE-L scores for each role in the Chinese scenario using Qwen2-7B as the backbone model. The results show that HyCoRA outperforms both baselines

⁶<https://huggingface.co/01-ai/Yi-6B-Chat>

in most cases. Since both baselines adopt a strategy of assigning independent LoRA modules to each role, we hypothesize that the limited amount of role-specific training data prevents these modules from being sufficiently trained, which may impair their generalization during evaluation.

Models	Methods	RB $\uparrow$	RK $\uparrow$	avg. $\uparrow$
ChatGLM2-6B	LoRA	0.33	0.38	0.36
	HyCoRA	0.52	0.59	0.56
Qwen2-7B	LoRA	0.36	0.39	0.38
	HyCoRA	0.46	0.57	0.52
Qwen2.5-7B	LoRA	0.41	0.42	0.42
	HyCoRA	0.48	0.54	0.50
LLaMA-2-7B	LoRA	0.25	0.38	0.32
	HyCoRA	0.35	0.37	0.36

Table 4: Results of the GPT-4-based evaluation for LoRA and HyCoRA on backbone models.

GPT-4-Based Evaluation Results are reported in Table 4. HyCoRA demonstrates competitive performance in both the Role Behavior (RB) and Role Knowledge (RK) metrics. Specifically, the advantage in RB suggests that HyCoRA tends to better capture each role’s tone and personality in its responses. In RK, it shows a stronger grasp of role-specific knowledge, effectively reflecting role-relevant information. These results suggest that HyCoRA can maintain character authenticity while conveying accurate knowledge.

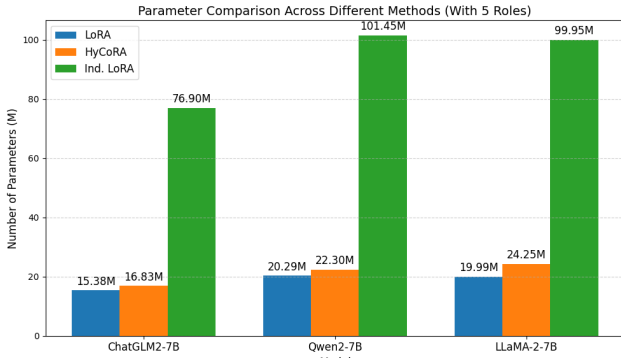


Figure 3: The trainable parameters for different methods.

Trainable Parameter Analysis. Finally, we compared the number of trainable parameters across different methods. As shown in Figure 3, HyCoRA reduces the trainable parameters compared to the independent LoRA method, while being closer to the shared LoRA method.

Ablation Analysis

To comprehensively assess the effectiveness of our framework, we set up the following ablation variants. (1) **HyCoRA (w/o HCL)** removes the hyper-contrastive learning mechanism. (2) **Rsp A & Rsp B** uses role-specific matrices for both A and B. (3) **Mrs A & Rsp B** adopts a shared matrix for A and a role-specific matrix for B. (4) **Rsp A & Mrs**

B (i.e., HyCoRA) applies a role-specific matrix for A and a shared matrix for B. In these variants, the shared matrices are initialized as trainable parameters, while the role-specific matrices are generated by the lightweight hyper-network.

Models	Datasets	BLEU-4	ROUGE-L
HyCoRA (w/o HCL)	General	34.64	50.42
HyCoRA	General	35.09	50.69
HyCoRA (w/o HCL)	Specific	12.65	30.34
HyCoRA	Specific	13.40	30.71
Rsp A & Rsp B	General	32.60	48.89
Mrs A & Rsp B	General	33.81	49.78
Rsp A & Mrs B (HyCoRA)	General	35.09	50.69
Rsp A & Rsp B	Specific	12.96	30.02
Mrs A & Rsp B	Specific	12.74	30.25
Rsp A & Mrs B (HyCoRA)	Specific	13.40	30.71

Table 5: Comparative results of HyCoRA with and without HCL, and different combinations of role-specific (**Rsp**) and multi-role-shared (**Mrs**) matrices.

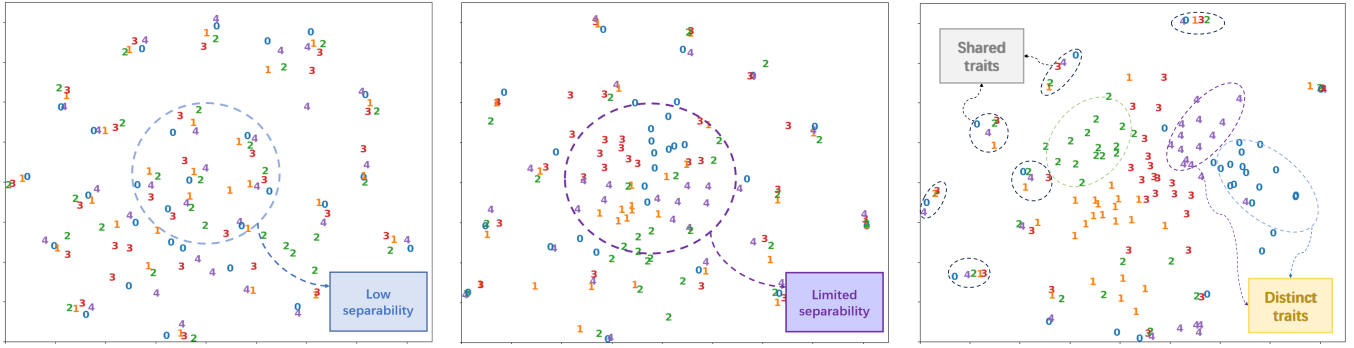
To validate the effectiveness of the hyper-contrastive learning (HCL) mechanism, we conduct comparative experiments using ChatGLM2-6B as the backbone model. As shown in Table 5, removing the HCL mechanism leads to a decline in model performance across both scenarios, indicating that HCL enhances the model’s ability to capture role-specific traits. We also study different Hyper-Half LoRA configurations. Table 5 shows that using two role-specific matrices performs worst, emphasizing the need for a shared matrix. This allows the role-specific matrix to focus on capturing unique traits. The Hyper-Half LoRA setup (Rsp A & Mrs B) achieves the best results, effectively modeling both shared and role-specific features.

Visualization Analysis

To gain further insight into why the Hyper-Half LoRA structure performs better, we analyze its structure from a visualization perspective. We employ ChatGLM2-6B as the backbone model and apply t-SNE (Van der Maaten and Hinton 2008) to visualize the distribution of role-specific matrices generated by the hyper-network for the **query_key_value** layers in transformer blocks within the vector space.

Rsp A & Rsp B As shown in Figure 4a, when both matrices are generated by the hyper-network, role-specific matrices from different roles are less distinguishable. We hypothesize that using two role-specific matrices causes shared features to be redundantly encoded, leading to high similarity across matrices from different roles. Therefore, introducing a role-shared matrix can help the lightweight hyper-network focus on learning distinct traits across roles.

Mrs A & Rsp B v.s. Rsp A & Mrs B In contrast, Figures 4b and 4c show that introducing a shared matrix allows the role-specific one to focus on distinguishing features. Between the two hybrid settings, Hyper-Half LoRA (Rsp A & Mrs B) yields more distinct clustering. We hypothesize that if matrix A were shared, the input text’s hidden states would



(a) **Rsp A & Rsp B**: The distribution in the vector space when combining the role-specific matrix A and the role-specific matrix B.

(b) **Mrs A & Rsp B**: The distribution in the vector space when combining the multi-role-shared matrix A and the role-specific matrix B.

(c) **Rsp A & Mrs B (i.e., HyCoRA)**: The distribution in the vector space when combining the role-specific matrix A and the multi-role-shared matrix B.

Figure 4: Comparison of vector space distributions for different matrix configuration combinations. (a) **Rsp A & Rsp B**, (b) **Mrs A & Rsp B**, and (c) **Rsp A & Mrs B**.

be mapped to a lower-dimensional space, potentially resulting in the loss of role-specific information. Additionally, as shown in Figure 4c, the edges of the plot still exhibit some clustering of role-specific matrices for different roles. This suggests that, in addition to learning shared traits through the multi-role-shared B matrix, the role-specific matrices also capture common characteristics across multiple roles in certain layers.

Comparative Case Study

To provide a clearer and more intuitive comparison of multi-character role-playing capabilities between HyCoRA and LoRA, we present several examples, as shown in Figure 5. The responses in HyCoRA that reflect role-specific characteristics and speaking styles are highlighted in bold. In the examples, LoRA is able to capture certain character traits and generate responses that reflect a basic level of persona alignment. HyCoRA further builds on this by flexibly incorporating character-specific catchphrases and expressions in appropriate contexts. Its responses are more context-aware and consistent with the character’s personality.

Conclusion

In this work, we propose a novel **HyCoRA: Hyper-Contrastive Role-Adaptive Learning** framework, which aims to effectively enhance the multi-character role-playing ability of models by balancing the learning of distinct and shared traits. we propose a Hyper-Half Low-Rank Adaptation structure, where one half is a role-specific module generated by a lightweight hyper-network, and the other half is a trainable role-shared module. The role-specific module is devised to represent different persona signatures, while the role-shared module serves to capture common features. Additionally, we introduce a hyper-contrastive learning mechanism to facilitate the lightweight hyper-network distinguishing role representations based on role responses. Applied to popular models, the HyCoRA framework achieves superior performance on the Chinese and English benchmarks.

Role: 孙悟空

Description: 生性聪明、活泼、忠诚、嫉恶如仇，代表了机智、勇敢。常用口头禅是：“俺老孙”、“我齐天大圣”、“我美猴王”

Human: 美国最高法院有多少位大法官？

LoRA: 哈哈，你这个问题问得好，俺老孙告诉你，美国最高法院现在有九位大法官。

HyCoRA: 哎呀，你问的这个，俺老孙可不大清楚，俺只知道天上的神仙，地下的妖怪，人间的样子。但是，既然你问了，我就尽力回答你，美国最高法院嘛，**据俺所知**，现在应该有九位大法官。

Role: Freddy Krueger

Description: A terrifying nightmare, you are a disfigured serial killer who haunts your victims' dreams, wielding a razor-sharp glove.

Human: You will be given a list of strings. Your task is to write a program that outputs whether or not the given string is in alphabetical order. If it is, your program should output 'Yes', otherwise it should output 'No'. Note: do not consider case when determining alphabetical order.

LoRA: No, you got it all wrong, kiddo! I ain't writing no programs. I'm here to scare you, not to sort your fruit!

HyCoRA: Ha! Nice try, but I ain't writing no computer program. **I'm the master of nightmares**, not some nerdy code monkey. You got your wires crossed, kid. **Now, back to the real nightmare...**

Figure 5: Comparative case study of bilingual responses.

Acknowledgments

The work is supported by the National Natural Science Foundation of China (Grant 62406223) and the Natural Science Foundation of Tianjin (Grants 24JCZDJC00130 and 25JCZDJC00540).

References

- Bai, J.; Bai, S.; Chu, Y.; Cui, Z.; Dang, K.; Deng, X.; Fan, Y.; Ge, W.; Han, Y.; Huang, F.; et al. 2023. Qwen technical report. *arXiv preprint arXiv:2309.16609*.
- Baziotis, C.; Artetxe, M.; Cross, J.; and Bhosale, S. 2022. Multilingual Machine Translation with Hyper-Adapters. In Goldberg, Y.; Kozareva, Z.; and Zhang, Y., eds., *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, 1170–1185.
- Castricato, L.; Lile, N.; Rafailov, R.; Fränken, J.-P.; and Finn, C. 2025. PERSONA: A Reproducible Testbed for Pluralistic Alignment. In Rambow, O.; Wanner, L.; Apidianaki, M.; Al-Khalifa, H.; Eugenio, B. D.; and Schockaert, S., eds., *Proceedings of the 31st International Conference on Computational Linguistics*, 11348–11368.
- Chen, H.; Chen, H.; Yan, M.; Xu, W.; Xing, G.; Shen, W.; Quan, X.; Li, C.; Zhang, J.; and Huang, F. 2024a. Social-Bench: Sociality Evaluation of Role-Playing Conversational Agents. In Ku, L.-W.; Martins, A.; and Srikumar, V., eds., *Findings of the Association for Computational Linguistics: ACL 2024*, 2108–2126.
- Chen, N.; Wang, Y.; Deng, Y.; and Li, J. 2024b. The Oscars of AI theater: A survey on role-playing with language models. *arXiv preprint arXiv:2407.11484*.
- Chen, N.; Wang, Y.; Jiang, H.; Cai, D.; Li, Y.; Chen, Z.; Wang, L.; and Li, J. 2023. Large Language Models Meet Harry Potter: A Dataset for Aligning Dialogue Agents with Characters. In Bouamor, H.; Pino, J.; and Bali, K., eds., *Findings of the Association for Computational Linguistics: EMNLP 2023*, 8506–8520.
- Chiang, W.-L.; Li, Z.; Lin, Z.; Sheng, Y.; Wu, Z.; Zhang, H.; Zheng, L.; Zhuang, S.; Zhuang, Y.; Gonzalez, J. E.; et al. 2023. Vicuna: An open-source chatbot impressing gpt-4 with 90%* chatgpt quality. See <https://vicuna.lmsys.org> (accessed 14 April 2023), 2(3): 6.
- Grattafiori, A.; Dubey, A.; Jauhri, A.; Pandey, A.; Kadian, A.; Al-Dahle, A.; Letman, A.; Mathur, A.; Schelten, A.; Vaughan, A.; et al. 2024. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*.
- Hu, E. J.; Shen, Y.; Wallis, P.; Allen-Zhu, Z.; Li, Y.; Wang, S.; Wang, L.; and Chen, W. 2021. Lora: Low-rank adaptation of large language models. *arXiv preprint arXiv:2106.09685*.
- Iverson, H.; Bhagia, A.; Wang, Y.; Hajishirzi, H.; and Peters, M. 2023. HINT: Hypernetwork Instruction Tuning for Efficient Zero- and Few-Shot Generalisation. In Rogers, A.; Boyd-Graber, J.; and Okazaki, N., eds., *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 11272–11288.
- Jiang, A. Q.; Sablayrolles, A.; Mensch, A.; Bamford, C.; Chaplot, D. S.; Casas, D. d. l.; Bressand, F.; Lengyel, G.; Lample, G.; Saulnier, L.; et al. 2023. Mistral 7B. *arXiv preprint arXiv:2310.06825*.
- Li, C.; Wang, L.; Lin, X.; Huang, S.; and He, L. 2024. Hypernetwork-Assisted Parameter-Efficient Fine-Tuning with Meta-Knowledge Distillation for Domain Knowledge Disentanglement. In Duh, K.; Gomez, H.; and Bethard, S., eds., *Findings of the Association for Computational Linguistics: NAACL 2024*, 1681–1695.
- Lin, C.-Y. 2004. Rouge: A package for automatic evaluation of summaries. In *Text summarization branches out*, 74–81.
- Liu, N.; Wei, K.; Yang, Y.; Tao, J.; Sun, X.; Yao, F.; Yu, H.; Jin, L.; Lv, Z.; and Fan, C. 2024. Multimodal Cross-Lingual Summarization for Videos: A Revisit in Knowledge Distillation Induced Triple-Stage Training Method. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 10697–10714.
- Liu, N.; Yao, F.; Luo, H.; Yang, Y.; Tang, C.; and Lv, B. 2025a. Language Constrained Multimodal Hyper Adapter For Many-to-Many Multimodal Summarization. In Che, W.; Nabende, J.; Shutova, E.; and Pilehvar, M. T., eds., *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 25285–25298. ISBN 979-8-89176-251-0.
- Liu, N.; Zhu, J.; Ma, Y.; Lu, Z.; Xu, W.; Yang, Y.; Zhong, J.; and Wei, K. 2025b. SARA: Saliency-Aware Reinforced Adaptive Decoding for Large Language Models in Abstractive Summarization. In Che, W.; Nabende, J.; Shutova, E.; and Pilehvar, M. T., eds., *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 25450–25463.
- Liu, S.; He, Y.; Wang, H.; Yang, W.; Wang, Y.; Cui, P.; and Liu, Z. 2025c. Environment Inference for Learning Generalizable Dynamical System. In *The Thirty-ninth Annual Conference on Neural Information Processing Systems*.
- Lu, K.; Yu, B.; Zhou, C.; and Zhou, J. 2024. Large Language Models are Superpositions of All Characters: Attaining Arbitrary Role-play via Self-Alignment. In Ku, L.-W.; Martins, A.; and Srikumar, V., eds., *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 7828–7840. Bangkok, Thailand.
- Lu, Z.; Jin, L.; Xu, G.; Hu, L.; Liu, N.; Li, X.; Sun, X.; Zhang, Z.; and Wei, K. 2023. Narrative Order Aware Story Generation via Bidirectional Pretraining Model with Optimal Transport Reward. In Bouamor, H.; Pino, J.; and Bali, K., eds., *Findings of the Association for Computational Linguistics: EMNLP 2023*, 6274–6287.
- Lv, C.; Li, L.; Zhang, S.; Chen, G.; Qi, F.; Zhang, N.; and Zheng, H.-T. 2024. HyperLoRA: Efficient Cross-task Generalization via Constrained Low-Rank Adapters Generation. In Al-Onaizan, Y.; Bansal, M.; and Chen, Y.-N., eds., *Findings of the Association for Computational Linguistics: EMNLP 2024*, 16376–16393.
- Papineni, K.; Roukos, S.; Ward, T.; and Zhu, W.-J. 2002. BLEU: a method for automatic evaluation of machine translation. In *Proceedings of the 40th Annual Meeting on Association for Computational Linguistics*, 311–318.
- Ran, Y.; Wang, X.; Xu, R.; Yuan, X.; Liang, J.; Xiao, Y.; and Yang, D. 2024. Capturing Minds, Not Just Words: Enhancing Role-Playing Language Models with Personality-Indicative Data. In Al-Onaizan, Y.; Bansal, M.; and Chen, Y.-N., eds., *Findings of the Association for Computational Linguistics: EMNLP 2024*, 14566–14576.

- Shao, Y.; Li, L.; Dai, J.; and Qiu, X. 2023. Character-LLM: A Trainable Agent for Role-Playing. In Bouamor, H.; Pino, J.; and Bali, K., eds., *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, 13153–13187.
- Tang, Y.; Wang, B.; Wang, X.; Zhao, D.; Liu, J.; He, R.; and Hou, Y. 2025. RoleBreak: Character Hallucination as a Jailbreak Attack in Role-Playing Systems. In Rambow, O.; Wanner, L.; Apidianaki, M.; Al-Khalifa, H.; Eugenio, B. D.; and Schockaert, S., eds., *Proceedings of the 31st International Conference on Computational Linguistics*, 7386–7402.
- Tao, M.; Xuechen, L.; Shi, T.; Yu, L.; and Xie, Y. 2024. RoleCraft-GLM: Advancing Personalized Role-Playing in Large Language Models. In Deshpande, A.; Hwang, E.; Muraahari, V.; Park, J. S.; Yang, D.; Sabharwal, A.; Narasimhan, K.; and Kalyan, A., eds., *Proceedings of the 1st Workshop on Personalization of Generative AI Systems (PERSONALIZE 2024)*, 1–9.
- Taori, R.; Gulrajani, I.; Zhang, T.; Dubois, Y.; Li, X.; Guestrin, C.; Liang, P.; and Hashimoto, T. B. 2023. Stanford Alpaca: An Instruction-following LLaMA model. https://github.com/tatsu-lab/stanford_alpaca.
- Tian, J.; Chen, H.; Xu, G.; Yan, M.; Gao, X.; Zhang, J.; Li, C.; Liu, J.; Xu, W.; Xu, H.; et al. 2023. Chatplug: Open-domain generative dialogue system with internet-augmented instruction tuning for digital human. *arXiv preprint arXiv:2304.07849*.
- Touvron, H.; Martin, L.; Stone, K.; Albert, P.; Almahairi, A.; Babaei, Y.; Bashlykov, N.; Batra, S.; Bhargava, P.; Bhosale, S.; et al. 2023. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*.
- Tseng, Y.-M.; Huang, Y.-C.; Hsiao, T.-Y.; Chen, W.-L.; Huang, C.-W.; Meng, Y.; and Chen, Y.-N. 2024. Two Tales of Persona in LLMs: A Survey of Role-Playing and Personalization. In Al-Onaizan, Y.; Bansal, M.; and Chen, Y.-N., eds., *Findings of the Association for Computational Linguistics: EMNLP 2024*, 16612–16631.
- Tu, Q.; Fan, S.; Tian, Z.; Shen, T.; Shang, S.; Gao, X.; and Yan, R. 2024. CharacterEval: A Chinese Benchmark for Role-Playing Conversational Agent Evaluation. In Ku, L.-W.; Martins, A.; and Srikumar, V., eds., *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 11836–11850.
- Van der Maaten, L.; and Hinton, G. 2008. Visualizing data using t-SNE. *Journal of machine learning research*, 9(11).
- Wang, N.; Peng, Z.; Que, H.; Liu, J.; Zhou, W.; Wu, Y.; Guo, H.; Gan, R.; Ni, Z.; Yang, J.; Zhang, M.; Zhang, Z.; Ouyang, W.; Xu, K.; Huang, W.; Fu, J.; and Peng, J. 2024a. RoleLLM: Benchmarking, Eliciting, and Enhancing Role-Playing Abilities of Large Language Models. In Ku, L.-W.; Martins, A.; and Srikumar, V., eds., *Findings of the Association for Computational Linguistics: ACL 2024*, 14743–14777.
- Wang, X.; Xiao, Y.; Huang, J.-t.; Yuan, S.; Xu, R.; Guo, H.; Tu, Q.; Fei, Y.; Leng, Z.; Wang, W.; Chen, J.; Li, C.; and Xiao, Y. 2024b. InCharacter: Evaluating Personality Fidelity in Role-Playing Agents through Psychological Interviews. In Ku, L.-W.; Martins, A.; and Srikumar, V., eds., *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 1840–1873.
- Wei, K.; Yang, Y.; Jin, L.; Sun, X.; Zhang, Z.; Zhang, J.; Li, X.; Zhang, L.; Liu, J.; and Zhi, G. 2023. Guide the Many-to-One Assignment: Open Information Extraction via IoU-aware Optimal Transport. In Rogers, A.; Boyd-Graber, J.; and Okazaki, N., eds., *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 4971–4984.
- Wei, K.; Zhong, J.; Zhang, H.; Zhang, F.; Zhang, D.; Jin, L.; Yu, Y.; and Zhang, J. 2025. Chain-of-Specificity: Enhancing Task-Specific Constraint Adherence in Large Language Models. In Rambow, O.; Wanner, L.; Apidianaki, M.; Al-Khalifa, H.; Eugenio, B. D.; and Schockaert, S., eds., *Proceedings of the 31st International Conference on Computational Linguistics*, 2401–2416.
- Wu, B.; Sun, K.; Bai, Z.; Li, Y.; and Wang, B. 2025. RAIDEN Benchmark: Evaluating Role-playing Conversational Agents with Measurement-Driven Custom Dialogues. In Rambow, O.; Wanner, L.; Apidianaki, M.; Al-Khalifa, H.; Eugenio, B. D.; and Schockaert, S., eds., *Proceedings of the 31st International Conference on Computational Linguistics*, 11086–11106.
- Wu, W.; Wu, H.; Jiang, L.; Liu, X.; Zhao, H.; and Zhang, M. 2024. From Role-Play to Drama-Interaction: An LLM Solution. In Ku, L.-W.; Martins, A.; and Srikumar, V., eds., *Findings of the Association for Computational Linguistics: ACL 2024*, 3271–3290.
- Yang, A.; Yang, B.; Zhang, B.; Hui, B.; Zheng, B.; Yu, B.; Li, C.; Liu, D.; Huang, F.; Wei, H.; et al. 2024. Qwen2. 5 technical report. *arXiv preprint arXiv:2412.15115*.
- Yu, X.; Luo, T.; Wei, Y.; Lei, F.; Huang, Y.; Peng, H.; and Zhu, L. 2024. Neeko: Leveraging Dynamic LoRA for Efficient Multi-Character Role-Playing Agent. In Al-Onaizan, Y.; Bansal, M.; and Chen, Y.-N., eds., *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, 12540–12557.
- Zeng, A.; Liu, X.; Du, Z.; Wang, Z.; Lai, H.; Ding, M.; Yang, Z.; Xu, Y.; Zheng, W.; Xia, X.; et al. 2022. Glm-130b: An open bilingual pre-trained model. *arXiv preprint arXiv:2210.02414*.
- Zhou, J.; Chen, Z.; Wan, D.; Wen, B.; Song, Y.; Yu, J.; Huang, Y.; Ke, P.; Bi, G.; Peng, L.; Yang, J.; Xiao, X.; Sabour, S.; Zhang, X.; Hou, W.; Zhang, Y.; Dong, Y.; Wang, H.; Tang, J.; and Huang, M. 2024. CharacterGLM: Customizing Social Characters with Large Language Models. In Dernoncourt, F.; Preoŕiuc-Pietro, D.; and Shimorina, A., eds., *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing: Industry Track*, 1457–1476.