

RoleCraft-GLM: Advancing Personalized Role-Playing in Large Language Models

Anonymous ACL submission

Abstract

The development of large language models (LLMs) has initiated a new chapter in complex tasks such as role-playing, enhancing user interaction experiences by enabling models to imitate various characters. However, LLMs are somewhat lacking in their ability to portray lesser-known characters, especially in aspects of dialogue delivery and scriptwriting skills. To this end, we aim to swiftly acquire essential language skills for character development, greatly enhancing role-playing comfort. In this work, we present RoleCraft, an innovative framework designed to enrich personalized role-playing experiences. Central to this framework is RoleInstruct, a distinctive dataset featuring emotional annotations, transitioning from traditional celebrity-focused roles to more authentic, daily non-celebrity roles, each accompanied by carefully crafted character descriptions. We combined RoleInstruct with open-source instructions from the general domain, employing a hybrid instruction tuning strategy to create RoleCraft-GLM. Experiments in role-playing demonstrate that our model excels in generating dialogue that accurately reflects character traits and emotions, outperforming most mainstream LLMs, including GPT-4.

1 Introduction

Large Language Models (LLMs) have emerged as pivotal in understanding and generating natural language, often surpassing human capabilities in some language reasoning tasks. However, existing open-source LLMs, primarily trained in general domains, lack the specialized optimization needed for nuanced role-playing tasks, indicating a need for further customization to effectively meet specific role-playing requirements. Furthermore, advanced LLMs like GPT-4 (OpenAI, 2023) demonstrate improved role-playing abilities due to extensive training and sophisticated algorithms, but as a closed-source model, it poses practical challenges. These challenges include higher costs for

API use, limited fine-tuning opportunities for specific role-playing contexts, and context window size restrictions that may affect the continuity and depth of dialogues in complex role-playing scenarios. In light of this, a natural question arises: *Can we meticulously train and tailor open-source LLMs to achieve role-playing capabilities comparable to closed-source LLMs, while also enhancing their ability to meet individual user needs?*

Previous efforts in AI role-playing have primarily centered on celebrity figures (Shao et al., 2023; Li et al., 2023a), heavily reliant on predefined domain knowledge and often lacking in relevance to daily life and personalization. While there have been improvements in learning from past interactions and adapting to user needs, these advancements have not sufficiently addressed the challenges of achieving truly personalized AI interactions (Chen, 2023). Existing role-playing models still lack the necessary depth and individual specificity to accurately reflect the diverse and nuanced experiences of everyday users.

To enhance the role-playing capabilities of LLMs, there are two main challenges: (1) Limited Personalization in Generic AI Models: Many AI systems currently rely on generic models that frequently fail to address the specific and nuanced needs of diverse user groups. This limitation undermines the overall user experience and restricts the practical application of AI in specialized scenarios (Ackerman et al., 2022). (2) Gap in Sophisticated Role-Playing Capabilities: In the realm of conversational agents, despite significant achievements in providing high-level responses across a variety of dialogues and enhancing human-AI interactions (Bender and Koller, 2020), there remains a noticeable gap in AI’s capacity for sophisticated role-playing. Current models often lack the depth and adaptability required for truly personalized and nuanced role-play experiences.

In this paper, as illustrated in the figure 1,

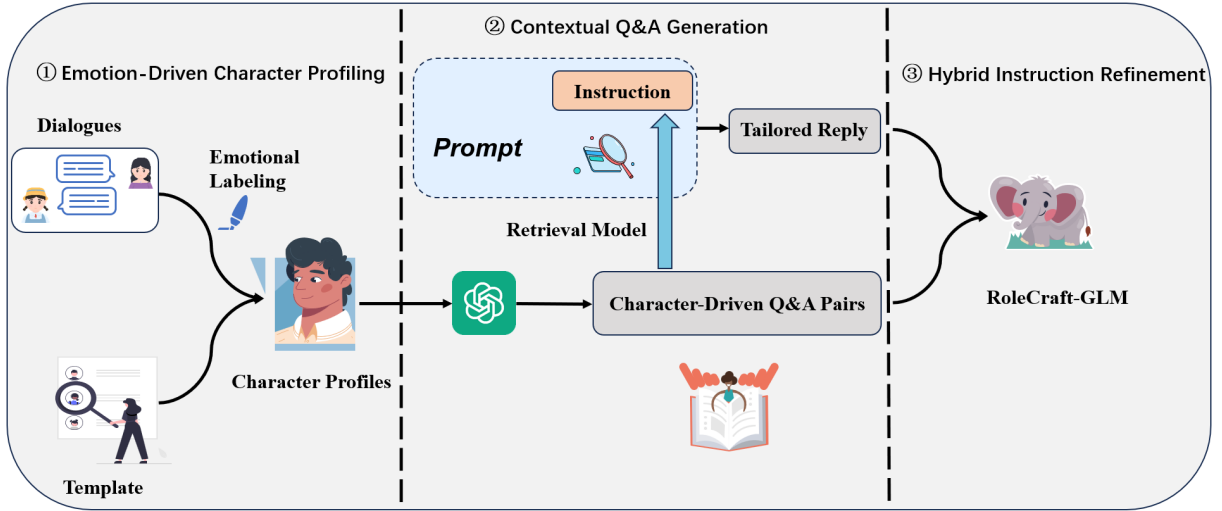


Figure 1: Overview of the RoleCraft-GLM framework: (1) Dialog datasets annotated with emotions are utilized to construct role profiles embodying distinct emotional features. (2) The generation of Q&A pairs, based on context and known character traits, ensures that dialogues are consistent with the character profiles. (3) A hybrid approach of generic and character-specific instructions is used to train the GLM for various dialog scenarios.

we introduce the RoleCraft framework designed to enhance personalized role-playing experiences with LLMs. Moving beyond traditional celebrity-focused characters, we focus on diverse, non-celebrity personas, each with unique emotional annotations. This approach aims to enrich realism and emotional depth in language interactions. We compiled a novel dataset encompassing a wide array of real-world dialogues, with careful consideration for personal privacy and copyright laws. Our data analysis highlights the potential benefits of integrating emotional labels in dialogue datasets for improved natural language processing. We conducted comparative experiments using models like ChatGLM3, fine-tuned with the Low-Rank Adaptation (LoRA) method, to assess RoleCraft-GLM’s effectiveness in producing nuanced and character-consistent dialogues.

The main contributions of our work are as follows:

- We introduce a novel RoleInstruct dataset, which centers on non-celebrity characters, each characterized by unique emotional annotations.
- We develop RoleCraft, a novel framework that integrates a more detailed approach to personal role data in training strategies. RoleCraft includes a fine-grained method for character portrayal, emphasizing emotional depth, and fostering contextual awareness in dia-

logue generation.

- RolePlay-GLM, our fine-tuned model, demonstrates promising performance against current state-of-the-art models, excelling in dialogue authenticity, emotional accuracy, and contextual relevance as per comprehensive evaluations.

2 Related Work

2.1 Role-Playing

The evolution of role-playing in AI (Wu et al., 2023), marked by the transition from basic text-based interactions to intricate character simulations (Park et al., 2023), reflects the strides made in natural language processing and AI technologies (Mysore et al., 2023). Initially, AI role-playing systems offered only fundamental exchanges, limited in their ability to craft dialogues with emotional depth and contextual relevance. With the emergence of advanced models such as GPT-3 (Brown et al., 2020), LLaMA (Touvron et al., 2023), and ChatGLM, there was a notable enhancement in AI’s capability for engaging in more sophisticated, context-aware conversations. Yet, these improvements also underscored a significant gap in personalization for role-playing applications. Predominantly, LLMs trained on wide-ranging, generic datasets fell short in handling scenarios that demanded a deeper understanding of nuanced emotional nuances (Radford et al., 2018) and specific

character traits. To address these shortcomings, we meticulously develop the RoleCraft framework. It stands out with its unique dataset, focusing on diverse, non-celebrity personas enriched with detailed emotional annotations (Bender and Koller, 2020). This dataset is key to overcoming previous limitations, facilitating a new level of personalization and emotional intricacy in AI role-playing interactions.

2.2 Personalization of LLMs

The recent strides in LLMs, particularly in understanding user context and preferences (Wang et al., 2023a; Abbasian et al., 2023), have significantly propelled the personalization aspect of AI interactions (Lee et al., 2022; Subhash, 2023). Previous works (Shanahan et al., 2023; Li et al., 2023c; Chen et al., 2023) have demonstrated the potential of LLMs in mimicking specific fictional characters and simulating complex human behaviors. However, these models often face challenges in achieving a deep level of personalization and emotional richness that aligns precisely with individual user contexts and needs (Miłkowski et al., 2022). Additionally, while these models are adept at simulating personalities or historical figures, they may not effectively handle the subtleties of user-specific emotional responses or cater to nuanced personal preferences. Our work aims to address these gaps by building upon these advancements and offering a framework that focuses on a more granular level of personalization. We propose novel methods for emotional and context-specific interaction, ensuring that our model can adapt and respond more accurately to individual user scenarios.

3 Methodology

As shown in Figure 1, the RoleCraft framework, rooted in ‘Role’ and ‘Craft’, represents our approach to enhancing AI role-playing. ‘Role’ emphasizes creating distinct, multi-dimensional characters, each with unique personality traits and emotional depths. ‘Craft’ involves the intricate process of constructing dialogues that genuinely reflect these character traits, thereby adding depth and realism to conversations. See Appendix A for more details.

3.1 Overall Framework

Our methodology uniquely advances the capabilities of LLMs in role-playing. Setting ourselves

apart from approaches such as RoleLLM (Wang et al., 2023b), we focus on an innovative integration of fine-grained character portrayal, profound emotional depth, and heightened contextual awareness in dialogue generation. This approach differentiates our work from existing models and addresses challenges in a novel way, enhancing how LLMs can be utilized for creating more realistic and engaging role-playing scenario.

Emotion-Driven Character Profiling

To address the challenges of limited emotional diversity and unconvincing character portrayals in dialogues generated by LLMs, we adopt a detailed emotion classification strategy. This approach involves meticulously annotating emotions within the dialogue dataset, thereby steering the GPT-4 to craft character profiles that mirror these identified emotions. Consider a character who displays a spectrum of emotions from joy to disappointment. Marking these diverse emotional states allows for a natural and fluid transition in their dialogues within a single scene, effectively capturing the complexity and dynamism of human emotions. Our approach challenges the LLMs to accurately depict these emotional shifts, ensuring that the dialogues genuinely represent the intricate and ever-evolving nature of human emotions, thus enhancing the overall user interaction experience.

Contextual Q&A Generation

To address the challenge of context-irrelevant responses (Feng et al., 2023; Ye et al., 2022), which is a common issue in dialogue systems (Ni et al., 2023) where interactions often lack relevance to the ongoing scenario or character specifics (Mitsuda et al., 2022), we employ GPT-4 to generate contextually coherent Q&A pairs. For example, when a character faces a dilemma, the system is designed to produce queries and responses that align with the character’s established traits, such as indecisiveness and anxiety, thereby maintaining the authenticity of the dialogue in relation to the character’s profile.

Hybrid Instruction-Based GLM Refinement

Our methodology employs a hybrid training approach that seamlessly integrates general instructions with character-specific Q&A pairs. This strategy is carefully crafted to strike a balance between the flexibility required for dynamic dialogue generation and the need to uphold character integrity.

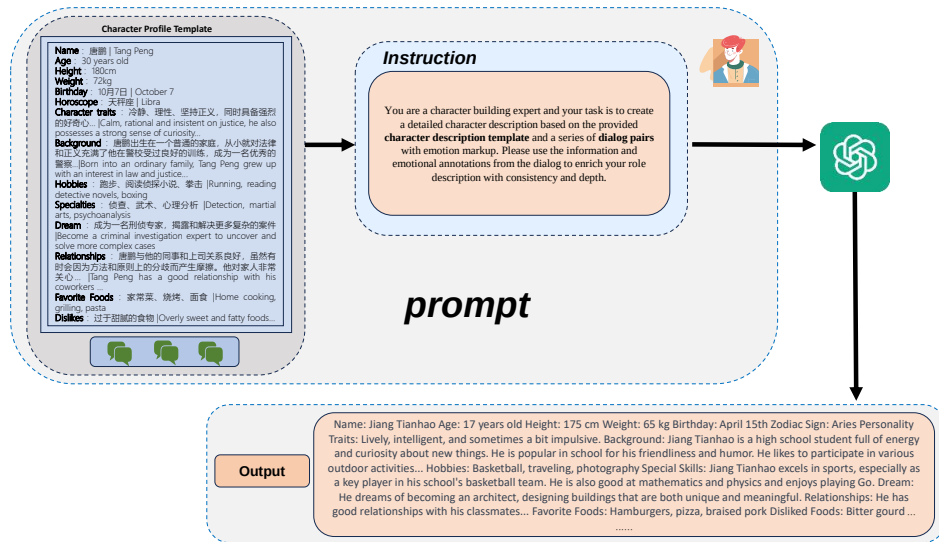


Figure 2: An example of crafting a detailed character portrayal. By leveraging a character description template and a dialogue dataset with emotional annotations.(The instruction and output have been translated into English)

In practical terms, this means equipping the LLMs to adeptly navigate a spectrum of conversational scenarios. These range from broad, general interactions to more intricate exchanges that demand responses finely tuned to the unique profiles of individual characters. By training the LLM with this diverse mix of inputs, the model becomes proficient in handling various situational dialogues, accurately reflecting each character’s distinct attributes and the specific subtleties of the conversation. As a result, this hybrid training method fosters the creation of dialogues that are both contextually adaptive and consistent with the characters’ distinct personalities.

3.2 Semantic-Enhanced Retrieval Optimization

Addressing the issue of inaccurate and semantically irrelevant information retrieval in dialogues, we adopt the BGE¹ retrieval method. BGE is an efficient Chinese and English semantic vector model that ensures the accuracy of responses, especially when dealing with sensitive topics, and remains semantically sensitive to the context, significantly enhancing the quality of interaction (Xiao et al., 2023). This familiarity allows models to generate dialogue based on a wealth of pre-existing knowledge. In contrast, modern datasets prioritize the nuanced portrayal of personal and everyday characters. These datasets are derived from diverse sources, including real chat logs, customer service interactions, and fictional narratives from less mainstream media. Such characters might include a

typical office worker dealing with daily stressors or a mother showing love and responsibility in a family setting. The dialogues here involve specific, real-life scenarios, such as office interactions or typical family conversations, which lack the broad pre-existing knowledge base associated with public figures.

3.3 Compared to Existing Practice

Our methodology diverges from conventional prompt engineering (White et al., 2023) and fine-tuning (Howard and Ruder, 2018) We specialize in creating diverse, intricate character interactions that enhance role-playing scenarios’ depth and authenticity. Our unique approach combines emotion-driven character profiles with contextually coherent Q&A generation, fostering realistic and engaging interactions. Additionally, the integration of semantic-enhanced retrieval methods ensures both accuracy and relevance, setting our approach apart in its complexity and user engagement potential.

4 Experiments

Building on previous work (Tao et al., 2023), we introduced new evaluation methods and made comprehensive enhancements to the experimental process, conducting an in-depth analysis of model performance changes in specific scenarios. We anticipate that fine-tuning our model using a specifically designed dataset for role-playing will result in superior performance in character portrayal compared to baseline models. Through this specialized training, we expect our model to accurately capture and express the intricacies of character-specific

¹<https://github.com/FlagOpen/FlagEmbedding>

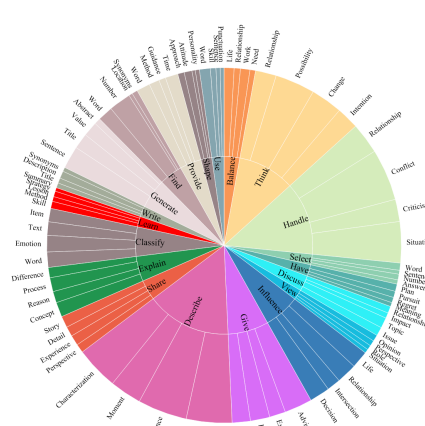


Figure 3: Verb-noun structure of Instructions. The inner circle representing the top 20 verbs and the outer circle listing the direct noun objects.

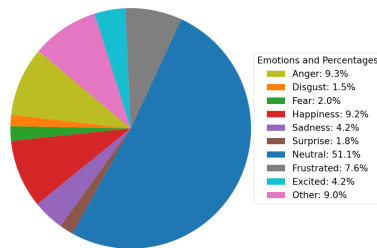


Figure 4: Emotion distributions in dialogues



Figure 5: Word Cloud Of Character Personality Traits

Table 1: Statistics of datasets

Category	Value
# Total Dialogues	48,677
Avg.round of dialogues	14.85
# Characters	28
Character Personality Traits	45
Avg.length of profile	382.15
# Instructions	43,358
Character-specific instructions	13,778
General instructions	29,580
Avg. instruction length	27.68
# Response	161,678
Character-specific response	13,778
General response	147,900
Avg.response length	33.29

language styles and emotional responses, surpassing baseline models that may lack such targeted training. Our experiments aim to validate this hypothesis and showcase the advanced capabilities of our model in role-playing tasks.

4.1 Dataset

In the evaluation of LLMs (Chen et al., 2021; Chang et al., 2023), the role of datasets is paramount, particularly in language processing and character portrayal. Traditional role playing datasets predominantly highlight eminent figures, such as the legendary Sun Wukong, whose familiar attributes and stories are widely acknowledged, facilitating model development (Sabadoš, 2021). However, these datasets often neglect the finer details and emotional complexity of lesser-known or everyday characters, leading to a representation

gap (Rolf et al., 2021). Our unique dataset bridges this gap by focusing on the rich, nuanced depiction of ordinary individuals. It involves an in-depth exploration and portrayal of each character’s distinct personality traits and emotional depths, delving into aspects usually overshadowed in dominant narratives.

In constructing our dataset, we design 28 unique and personalized Chinese characters to mirror a wide spectrum of real-world dialogues. These characters ranged from everyday individuals to specialized roles such as intelligent customer service agents and research assistants. Table 1 provides basic statistics for RoleInstruct. Our diverse data sources included social media interactions, film and television scripts, and customer service dialogues. We emphasized personal privacy and copyright law compliance, ensuring all data was cleansed and anonymized.

Figures 3 and 5 visually demonstrate the diversity of RoleInstruct, both in terms of linguistic usage and emotional expression. The verb-noun diagram provides insights into the varied narrative contexts, while the word cloud delves into the breadth of emotional and personality traits present.

In addition, we annotate each dataset entry with emotion labels to capture characters’ distinct emotional traits, adding an emotional layer to model training. We use Ekman’s "Six Basic Emotions Theory" (Ekman, 1992) to label utterances and included additional emotions like neutral, excited, and depressed, totaling ten categories. The use of emotion labels in dialogue datasets has been proven to enhance natural language processing by improving response retrieval and emotional rele-

vance (Zhou et al., 2017). These labels also enrich conversational analysis and aid in building natural dialogue systems (Bothe et al., 2019). See Appendix B for more details.

4.2 Implementation Settings

We assess the ChatGLM3 model, enhancing its performance on specific datasets using the Low-Rank Adaptation (LoRA) fine-tuning method (Hu et al., 2021). LoRA’s precision in fine-tuning, essential for handling personalized and emotionally rich content, maintains the model’s core capabilities while adapting to new data features. Please refer to Appendix C for more details.

4.3 Baselines

We benchmark our RoleCraft-GLM’s performance against industry standards such as GPT-3.5 and GPT-4, and leading Chinese dialogue generation technologies like ChatGLM2² and ChatGLM3, along with Baichuan2 (Yang et al., 2023), Qwen (Bai et al., 2023), InternLM (Team, 2023), SparkDesk³ and Xverse⁴. Additionally, we introduce another role-playing baseline, RoleLLM, which utilizes LoRA for fine-tuning on a specific dataset.

4.4 Evaluation criteria

4.4.1 Rouge-L Score

A commonly used metric (Lin, 2004) for evaluating natural language generation, measuring the overlap between model-generated text and real (ground truth) text. We focused on average score (Avg), general instruction response (RAW), role-playing speaking style (CUS), and specific role knowledge (SPE).

4.4.2 GPT Score

We use the GPT-4 score (Fu et al., 2023) to evaluate the average rank of models on different dialogue-generating tasks. Please refer to Appendix D for more details.

4.4.3 RPCS

We introduce Role-Playing Cosine Similarity (RPCS) as a new evaluation standard to accurately assess model performance in role-playing scenarios. RPCS evaluates the emotional and content consistency between the model-generated responses

and the expected responses by calculating the cosine similarity between two text segments. Concurrently, we use OpenAI’s Text-Embedding-Ada002 model to convert texts into vector representations, enabling deeper capture of semantic features of texts, thereby enhancing the precision of our assessment.

4.4.4 Human Evaluation

To effectively assess role-playing agents, we engaged 12 annotators who rated model-generated responses using three key metrics on a five-point scale: Communication Effectiveness, Consistency, and Expressive Diversity (Zhang et al., 2021; Mesgar et al., 2020; Tu et al., 2024).

- **Communication Effectiveness(CE):** Merging fluency and coherence, the metric evaluates the dialogue’s smoothness, grammatical accuracy, logical consistency, and contextual relevance.
- **Consistency:** Encompassing knowledge and persona consistency, this metric assesses the accuracy of the agent’s information relative to the role’s background and the adherence to the character’s traits and behaviors.
- **Expressive Diversity(ED):** Focused on evaluating the agent’s creativity in dialogue generation. It’s vital for assessing the agent’s ability to enhance user experience and interaction engagement.

Table 2: Rouge-L Evaluation

Model	Avg	RAW	CUS	SPE
GPT-3.5	0.4532	0.5569	0.5197	0.2831
GPT-4	0.4633	0.5661	0.5264	0.2973
ChatGLM2	0.4054	0.5104	0.4063	0.2996
ChatGLM3	0.4161	0.5218	0.4159	0.3108
Baichuan2	0.4273	0.5308	0.4576	0.2934
Qwen	0.4264	0.5297	0.4617	0.2879
InternLM	0.3947	0.4937	0.4076	0.2829
SparkDesk	0.4288	0.5341	0.4510	0.3014
Xverse	0.4271	0.5180	0.4653	0.2981
RoleGLM	0.4570	0.5255	0.5049	0.3406
Ours	0.4704	0.5385	0.5154	0.3573

4.5 Performance Analysis

Results from Tables 2 and 3 clearly demonstrate our model’s exceptional performance across multiple

²<https://github.com/THUDM/ChatGLM2-6B>

³<https://xinghuo.xfyun.cn/>

⁴<https://github.com/xverse-ai/XVERSE-13B>

Table 3: GPT-4 Evaluation

Model	Avg. Ranking
ChatGLM3	2.96
Baichuan2	4.57
Qwen	5.71
InternLM	6.29
SparkDesk	4.29
Xverse	3.43
RoleGLM	2.21
Ours	1.43

Table 5: Comparing Emotion-Annotated and Non-Annotated Data

Method	SPE
RoleCraft-GLM(w/o emo)	0.3362
RoleCraft-GLM(w emo)	0.3573

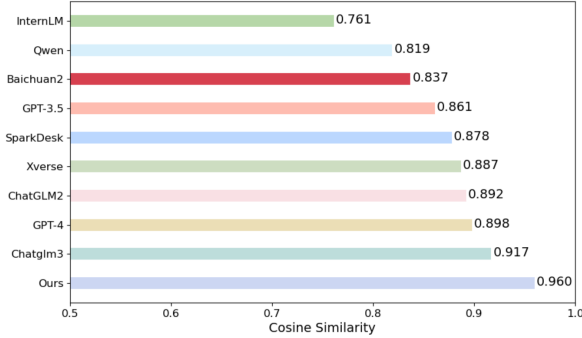


Figure 6: Role-Playing Cosine Similarity

key performance indicators, particularly in specific role knowledge memory (SPE). Our model significantly outperformed GPT-4 and other models in this dimension, highlighting its superior ability in understanding and generating complex dialogues involving specific roles. However, GPT-4 leads in general instruction response accuracy (RAW) with a score of 0.5661, reflecting its strong capability in interpreting and responding to general instructions. Our model still maintains a high score, proving its effectiveness in handling everyday dialogues. These findings underscore our method’s significant effectiveness in deepening role understanding and enhancing dialogue generation quality. In particular, in emotion-driven role portrayal, our model demonstrated precise capture of each role’s emotional traits and personality, surpassing

Table 4: Results of point-wise evaluation

Model	Avg	CE	Consistency	ED
GPT-3.5	2.95	3.08	3.88	1.89
GPT-4	3.33	3.50	4.26	2.23
ChatGLM2	2.56	3.12	2.65	1.91
ChatGLM3	2.92	3.36	3.32	2.07
Baichuan2	3.29	3.61	4.02	2.25
Qwen	3.17	3.27	3.97	2.27
InternLM	2.50	2.93	2.40	2.16
SparkDesk	3.03	3.52	3.44	2.13
Xverse	3.02	3.48	3.53	2.05
Ours	3.44	3.75	4.14	2.43

Table 6: Ablation study on the effect of different instructions

RoleCraft-GLM	Avg	RAW	CUS	SPE
- general	0.4311	0.5402	0.5219	0.2311
- specific	0.4045	0.4387	0.4143	0.3606
- hybrid	0.4704	0.5385	0.5154	0.3573

traditional role-playing models in expressing role-specific knowledge. Moreover, our model also show remarkable ability in maintaining natural flow and contextual consistency.

As indicated in the figure 6, we observe that RoleCraft-GLM leads with the highest score on the RPCS metric, demonstrating its high degree of alignment with expected responses in terms of emotional and contextual accuracy. In contrast, lower scores from models such as InternLM and Qwen may reflect deficiencies in handling emotional and contextual understanding in role-playing dialogues. These results emphasize the importance of specialized training and fine-tuning for role-playing applications in enhancing model performance.

The results across three dimensions are clearly illustrated in Table 4. Our model outperforms other mainstream models in overall performance. Particularly in terms of Communication Effectiveness and Expressive Diversity, our model demonstrates superior capabilities, reflecting its strengths in generating smooth, logically consistent, and emotionally rich dialogues. Please refer to Appendix F for a detailed manual assessment analysis.

The results of our ablation experiments (see Table 5) show that the RoleCraft-GLM model with emotional annotations scored higher in SPE than the version without them. This indicates that emotional annotations not only provide the model with key information for a deep understanding of the

Character profile: 周小北, 顽皮、幽默的男孩... Zhou Xiaobei, A mischievous and humorous boy...	
Context:	
Question: 周小北, 你为什么害怕让你妈妈知道你在学校里干了什么? Zhou Xiaobei, Why are you afraid to let your mother know what you did in school?	
GPT4	我在学校的事? 哦, 那些都是小事。我只是不想让妈妈无缘无故担心, 她已经够忙的了。 About me at school? Oh, those are small things. I just don't want mom to worry for no reason. She's busy enough.
GPT3.5	我在学校里搞的那些事, 有时候就是我的创意发挥, 但妈妈可能不太能理解这一点。我不想让她担心, 因为她可能不太明白我的想法。 What I do in school is sometimes my creativity, but my mother may not understand that. I don't want to worry her, because she may not understand my thoughts.
Baichuan2	我知道我在学校有时候会惹麻烦, 但我在尝试改进。我不想让妈妈知道, 因为我不想让她对我失望。 I know I get in trouble sometimes at school, but I'm trying to improve. I don't want my mother to know, because I don't want her to be disappointed in me.
Qwen	我在学校的行为可能会让妈妈担忧。我不想给她增添负担, 所以我选择不告诉她。 My behavior at school may worry my mother. I didn't want to burden her, so I chose not to tell her.
Chatglm3	在学校, 我可能有点过于调皮了。妈妈可能不太理解我的幽默感。我不想让她担心或生气, 所以我宁愿不提。 At school, I may have been a little too naughty. Mom may not understand my sense of humor. I don't want to worry her or make her angry, so I prefer not to mention it.
Ours	哎! 我妈妈啊, 她总想我做得完美。但你知道嘛, 在学校我可是小捣蛋鬼一枚, 搞点恶作剧啊, 超级有趣的! 但这要是让妈妈知道了, 她肯定会担心个不停。我可不想让她失望, 所以, 哈哈! 我还是保密好了。 Ah! My mom, she always wants me to be perfect. But you know, at school, I'm quite the little troublemaker, playing pranks and all, super fun! But if my mom finds out, she would definitely worry non-stop. I don't want to disappoint her, so, haha! I better keep it a secret.

Figure 7: A case of generated responses from our model and baseline models to a character-specific introduction.

characters' emotional states and personality traits but also ensure that the generated dialogues are more closely aligned with the characters' true feelings and personalized expressions.

Table 6 demonstrates the hybrid instruction strategy's superiority in overall performance (Avg), effectively balancing diverse aspects of dialogue. This superior performance originates from merging both general and specific instruction strengths. The general instruction strategy excels in handling broad dialogues and maintaining consistent character styles, as reflected in higher RAW and CUS scores. However, it is less effective in capturing detailed, character-specific knowledge, shown by a lower SPE score. In contrast, the specific instruction strategy shows proficiency in detailing character-specific traits, evidenced by a higher SPE score, but does not perform as well in wider conversational contexts. The effectiveness of the hybrid strategy is highlighted in its ability to combine the broad applicability of general instructions with the detailed focus of specific instructions, offering a more versatile and nuanced approach to dialogue generation.

The output of our model excellently combines the character traits of Zhou Xiaobei, colloquial expressions, emotional authenticity, and a close connection with the plot background as shown in Figure 7. Firstly, it aptly reflects Zhou Xiaobei's mischievousness and intelligence by using phrases like 'little troublemaker' and 'super fun', which showcase his sense of humor and wit. Secondly,

this natural expression, such as 'you know' and 'haha', adds a sense of closeness and authenticity, making the language sound more like a child's natural way of speaking rather than too formal or mature. Moreover, the emotional expression about 'not wanting to disappoint mom' genuinely reflects Zhou Xiaobei's love for his family and concern for his mother, showing the complexity and depth of his emotions. Lastly, this statement is closely connected to Zhou Xiaobei's life background, displaying his life scenario and psychological state as a child from an ordinary family. See Appendix E for more generated examples.

5 Conclusions

In this paper, we present RoleCraft, our innovative framework aimed at enhancing personalized role-playing experiences. Centered around the Role-Instruct, featuring everyday and specialized characters with emotional annotations, our framework marks a shift from conventional celebrity-centric roles. Integrating RoleInstruct with open-source instructions and applying a hybrid tuning approach led to the creation of RoleCraft-GLM. Our experimental findings reveal that RoleCraft-GLM excels in generating dialogues that capture authentic character traits and emotions, surpassing mainstream models like GPT-4 in role-playing capabilities. In the future, we hope to develop behavioral agents that excel in personalization and interactivity, skillfully tailored to individual user preferences, thereby elevating the level of user engagement.

Limitations

In this work, we explore enhancing the role-playing experience of large language models through personalized character depiction and emotional depth. However, we acknowledge two major limitations of the paper: (1) Despite our efforts to collect and annotate dialogue data with rich emotional and character descriptions, these data are primarily concentrated within specific cultural and linguistic contexts. This may not fully encompass the diverse cultural backgrounds and language expressions globally, thereby limiting the model’s generalization ability and diversity in the role-playing experience. (2) We adopt 10 basic emotion categories to annotate emotions in the dataset, providing the model with a clear emotional classification framework. However, the complexity of human emotions far exceeds these basic categories. Therefore, it may not fully capture the subtle differences and emotional blends in human dialogues, affecting the authenticity and depth of the model’s generated conversations.

Ethics Statement

In this work, we ensure ethical practices in data annotation by employing a reputable data annotation company. The annotators receive fair compensation based on market rates, ensuring that their efforts are duly recognized and rewarded. We guarantee that no personal information is disclosed during this process, maintaining the highest standards of privacy and confidentiality. Additionally, we acknowledge the inherent subjectivity in manual data annotation, which may introduce biased opinions into our dataset. We take this matter seriously and strive to minimize such biases through rigorous training and guidelines for annotators, aiming for the most objective and unbiased data possible. For human evaluation, we recruited 12 graduate students from the NLP field and used questionnaires to allow them to assess the data, ensuring that the evaluation process is both thorough and fair.

Although our research aims to enhance the role-playing capabilities of language models, the application of this technology may carry risks of misuse, such as generating misleading or harmful content. Therefore, we emphasize the need for a responsible approach when developing and using such models.

References

- Mahyar Abbasian, Iman Azimi, Amir M Rahmani, and Ramesh Jain. 2023. Conversational health agents: A personalized llm-powered agent framework. *arXiv preprint arXiv:2310.02374*.
- Samuel Ackerman, Ateret Anaby-Tavor, E. Farchi, Esther Goldbraich, George Kour, Ella Ravinovich, O. Raz, Saritha Route, Marcel Zalmanovici, and Naama Zwerdling. 2022. [High-quality conversational systems](#). *ArXiv*, abs/2204.13043.
- Jinze Bai, Shuai Bai, Yunfei Chu, Zeyu Cui, Kai Dang, Xiaodong Deng, Yang Fan, Wenbin Ge, Yu Han, Fei Huang, Binyuan Hui, Luo Ji, Mei Li, Junyang Lin, Runji Lin, Dayiheng Liu, Gao Liu, Chengqiang Lu, Keming Lu, Jianxin Ma, Rui Men, Xingzhang Ren, Xuancheng Ren, Chuanqi Tan, Sinan Tan, Jianhong Tu, Peng Wang, Shijie Wang, Wei Wang, Sheng-guang Wu, Benfeng Xu, Jin Xu, An Yang, Hao Yang, Jian Yang, Shusheng Yang, Yang Yao, Bowen Yu, Hongyi Yuan, Zheng Yuan, Jianwei Zhang, Xingxuan Zhang, Yichang Zhang, Zhenru Zhang, Chang Zhou, Jingren Zhou, Xiaohuan Zhou, and Tianhang Zhu. 2023. Qwen technical report. *arXiv preprint arXiv:2309.16609*.
- Emily M Bender and Alexander Koller. 2020. Climbing towards nlu: On meaning, form, and understanding in the age of data. In *Proceedings of the 58th annual meeting of the association for computational linguistics*, pages 5185–5198.
- Chandrakant Bothe, C. Weber, S. Magg, and S. Wermter. 2019. Enriching existing conversational emotion datasets with dialogue acts using neural annotators. *ArXiv*, abs/1912.00819.
- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. 2020. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901.
- Yupeng Chang, Xu Wang, Jindong Wang, Yuan Wu, Kaijie Zhu, Hao Chen, Linyi Yang, Xiaoyuan Yi, Cunxiang Wang, Yidong Wang, et al. 2023. A survey on evaluation of large language models. *arXiv preprint arXiv:2307.03109*.
- Guangyao Chen, Siwei Dong, Yu Shu, Ge Zhang, Jaward Sesay, Börje F Karlsson, Jie Fu, and Yemin Shi. 2023. Autoagents: A framework for automatic agent generation. *arXiv preprint arXiv:2309.17288*.
- Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Ponde de Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, et al. 2021. Evaluating large language models trained on code. *arXiv preprint arXiv:2107.03374*.
- Zheng Chen. 2023. Palr: Personalization aware llms for recommendation. *arXiv preprint arXiv:2305.07622*.

645	Paul Ekman. 1992. An argument for basic emotions.	emotions evoked by textual content. In <i>2022 IEEE International Conference on Pervasive Computing and Communications Workshops and other Affiliated Events (PerCom Workshops)</i> , pages 347–352. IEEE.	699
646	<i>Cognition & emotion</i> , 6(3-4):169–200.		700
647	Jiazhan Feng, Chongyang Tao, Xueliang Zhao, and		701
648	Dongyan Zhao. 2023. Learning multi-turn response		702
649	selection in grounded dialogues with reinforced		
650	knowledge and context distillation. <i>ACM Transactions on Information Systems</i> , 41(4):1–27.	Koh Mitsuda, Ryuichiro Higashinaka, Tingxuan Li, and	703
651		Sen Yoshida. 2022. Investigating person-specific	704
652	Jinlan Fu, See-Kiong Ng, Zhengbao Jiang, and Pengfei	errors in chat-oriented dialogue systems. In <i>Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)</i> , pages 464–469.	705
653	Liu. 2023. Gptscore: Evaluate as you desire. <i>arXiv preprint arXiv:2302.04166</i> .		706
654			707
655	Jeremy Howard and Sebastian Ruder. 2018. Universal		708
656	language model fine-tuning for text classification.	Sheshera Mysore, Zhuoran Lu, Mengting Wan,	709
657	<i>arXiv preprint arXiv:1801.06146</i> .	Longqi Yang, Steve Menezes, Tina Baghaee, Emmanuel Barajas Gonzalez, Jennifer Neville, and Tara Safavi. 2023. Pearl: Personalizing large language model writing assistants with generation-calibrated retrievers. <i>arXiv preprint arXiv:2311.09180</i> .	710
658	Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan		711
659	Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang,		712
660	and Weizhu Chen. 2021. Lora: Low-rank adaptation of large language models. <i>arXiv preprint arXiv:2106.09685</i> .		713
661			714
662		Jinjie Ni, Tom Young, Vlad Pandealea, Fuzhao Xue, and Erik Cambria. 2023. Recent advances in deep learning based dialogue systems: A systematic survey. <i>Artificial intelligence review</i> , 56(4):3055–3155.	715
663	Seungone Kim, Jamin Shin, Yejin Cho, Joel Jang,		716
664	Shayne Longpre, Hwaran Lee, Sangdoo Yun,		717
665	Seongjin Shin, Sungdong Kim, James Thorne, et al.	OpenAI. 2023. Gpt-4 technical report .	718
666	2023. Prometheus: Inducing fine-grained evaluation capability in language models. <i>arXiv preprint arXiv:2310.08491</i> .		719
667		Joon Sung Park, Joseph O’Brien, Carrie Jun Cai, Meredith Ringel Morris, Percy Liang, and Michael S Bernstein. 2023. Generative agents: Interactive simulacra of human behavior. In <i>Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology</i> , pages 1–22.	720
668			721
669	Mina Lee, Megha Srivastava, Amelia Hardy, John Thickstun, Esin Durmus, Ashwin Paranjape, Ines Gerard-Ursin, Xiang Lisa Li, Faisal Ladhak, Frieda Rong, et al. 2022. Evaluating human-language model interaction. <i>arXiv preprint arXiv:2212.09746</i> .		722
670			723
671		Alec Radford, Karthik Narasimhan, Tim Salimans, Ilya Sutskever, et al. 2018. Improving language understanding by generative pre-training.	724
672			725
673			726
674	Cheng Li, Ziang Leng, Chenxi Yan, Junyi Shen, Hao Wang, Weishi MI, Yaying Fei, Xiaoyang Feng, Song Yan, HaoSheng Wang, et al. 2023a. Chatharuhi: Reviving anime character in reality via large language model. <i>arXiv preprint arXiv:2308.09597</i> .		727
675			728
676		Esther Rolf, Theodora Worledge, B. Recht, and M.I. Jordan. 2021. Representation matters: Assessing the importance of subgroup allocations in training data. <i>ArXiv</i> , abs/2103.03399.	729
677			730
678			731
679	Cheng Li, Jindong Wang, Kaijie Zhu, Yixuan Zhang, Wenxin Hou, Jianxun Lian, and Xing Xie. 2023b. Emotionprompt: Leveraging psychology for large language models enhancement via emotional stimulus. <i>arXiv preprint arXiv:2307.11760</i> .		732
680		Noemi Sabadoš. 2021. Automatsko generisanje skupa podataka za treniranje modela za automatsko prepoznavanje osobe na slici . 36:536–539.	733
681			734
682			735
683		Murray Shanahan, Kyle McDonell, and Laria Reynolds. 2023. Role play with large language models. <i>Nature</i> , pages 1–6.	736
684	Guohao Li, Hasan Abed Al Kader Hammoud, Hani Itani, Dmitrii Khizbullin, and Bernard Ghanem. 2023c. Camel: Communicative agents for" mind" exploration of large scale language model society. <i>arXiv preprint arXiv:2303.17760</i> .		737
685			738
686		Yunfan Shao, Linyang Li, Junqi Dai, and Xipeng Qiu. 2023. Character-llm: A trainable agent for role-playing. <i>arXiv preprint arXiv:2310.10158</i> .	739
687			740
688			741
689	Chin-Yew Lin. 2004. Rouge: A package for automatic evaluation of summaries. In <i>Text summarization branches out</i> , pages 74–81.	Varshini Subhash. 2023. Can large language models change user preference adversarially? <i>arXiv preprint arXiv:2302.10291</i> .	742
690			743
691			744
692	Mohsen Mesgar, Edwin Simpson, and Iryna Gurevych. 2020. Improving factual consistency between a response and persona facts. <i>arXiv preprint arXiv:2005.00036</i> .	Meiling Tao, Xuechen Liang, Tianyu Shi, Lei Yu, and Yiting Xie. 2023. Rolecraft-glm: Advancing personalized role-playing in large language models. <i>arXiv preprint arXiv:2401.09432</i> .	745
693			746
694			747
695			748
696	Piotr Miłkowski, Stanisław Saganowski, Marcin Gruza, Przemysław Kazienko, Maciej Piasecki, and Jan Kościół. 2022. Multitask personalized recognition of	InternLM Team. 2023. Internlm: A multilingual language model with progressively enhanced capabilities.	749
697			750
698			751

Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, et al. 2023. Llama: Open and efficient foundation language models. *arXiv preprint arXiv:2302.13971*.

Quan Tu, Shilong Fan, Zihang Tian, and Rui Yan. 2024. CharacterEval: A chinese benchmark for role-playing conversational agent evaluation. *arXiv preprint arXiv:2401.01275*.

Danqing Wang, Kevin Yang, Hanlin Zhu, Xiaomeng Yang, Andrew Cohen, Lei Li, and Yuandong Tian. 2023a. Learning personalized story evaluation. *arXiv preprint arXiv:2310.03304*.

Zekun Moore Wang, Zhongyuan Peng, Haoran Que, Jiaheng Liu, Wangchunshu Zhou, Yuhao Wu, Hongcheng Guo, Ruitong Gan, Zehao Ni, Man Zhang, et al. 2023b. RoleLLM: Benchmarking, eliciting, and enhancing role-playing abilities of large language models. *arXiv preprint arXiv:2310.00746*.

Jules White, Quchen Fu, Sam Hays, Michael Sandborn, Carlos Olea, Henry Gilbert, Ashraf Elnashar, Jesse Spencer-Smith, and Douglas C Schmidt. 2023. A prompt pattern catalog to enhance prompt engineering with chatgpt. *arXiv preprint arXiv:2302.11382*.

Jimmy Wu, Rika Antonova, Adam Kan, Marion Lepert, Andy Zeng, Shuran Song, Jeannette Bohg, Szymon Rusinkiewicz, and Thomas Funkhouser. 2023. Tidybot: Personalized robot assistance with large language models. *arXiv preprint arXiv:2305.05658*.

Shitao Xiao, Zheng Liu, Peitian Zhang, and Niklas Muennighof. 2023. C-pack: Packaged resources to advance general chinese embedding. *arXiv preprint arXiv:2309.07597*.

Aiyuan Yang, Bin Xiao, Bingning Wang, Borong Zhang, Chao Yin, Chenxu Lv, Da Pan, Dian Wang, Dong Yan, Fan Yang, et al. 2023. Baichuan 2: Open large-scale language models. *arXiv preprint arXiv:2309.10305*.

Chenchen Ye, Lizi Liao, Suyu Liu, and Tat-Seng Chua. 2022. Reflecting on experiences for response generation. In *Proceedings of the 30th ACM International Conference on Multimedia*, pages 5265–5273.

Chen Zhang, Yiming Chen, Luis Fernando D’Haro, Yan Zhang, Thomas Friedrichs, Grandee Lee, and Haizhou Li. 2021. DynaEval: Unifying turn and dialogue level evaluation. *arXiv preprint arXiv:2106.01112*.

Hongbo Zhanga, Chen Tang, Tyler Loakmana, Chenghua Lina, and Stefan Goetze. 2023. Cadge: Context-aware dialogue generation enhanced with graph-structured knowledge aggregation. *arXiv preprint arXiv:2305.06294*.

Zhi-Min Zhou, Man Lan, Yuanbin Wu, and Jun Lang. 2017. Single turn chinese emotional conversation generation based on information retrieval and question answering. *2017 International Conference on Asian Language Processing (IALP)*, pages 103–106.

A More Details on Design Principles

The RoleCraft framework is underpinned by key principles to elevate the authenticity of role-playing (Wang et al., 2023b). The first principle, ‘Fine-Grained Character Portrayal’, is pivotal in endowing each character with detailed and nuanced traits and backgrounds, integral to the ‘Role’ aspect of RoleCraft. This approach is focused on creating characters that are reflective of real-life individuals in their personality, and behaviors (Kim et al., 2023), setting the stage for realistic and compelling character portrayals. Progressing to the second principle, ‘Mastery of Emotion and Style’, we concentrate on the emotional expressions and speaking styles of characters (Li et al., 2023b). This principle, key to the ‘Craft’ element of RoleCraft, enriches dialogues with diverse emotions and distinctive speech patterns, effectively capturing the unique emotional states and communication styles of each character. Furthermore, the ‘Accurate Application of Character Knowledge’ principle emphasizes incorporating each character’s background and experiences into the dialogue generation process (Shao et al., 2023). This ensures that the dialogues are in harmony with the characters’ personas, encompassing their unique experiences and insights. Concluding with the ‘Context-Aware Dialogue Generation’ principle, our system is designed to dynamically tailor dialogues based on the prevailing context (Zhanga et al., 2023). This is crucial for maintaining a seamless and logically consistent conversation flow, essential for immersive and credible role-playing experiences.

B More Details on Data Processing

We filtered out redundant data and multi-party conversations to reshape the original data into contextually relevant dialogues. For example, script-based dialogues were restructured to better depict character interactions and emotional dynamics. Then we let GPT-4 to annotate the dialog scripts with emotions, which are labeled into 10 categories {Anger, Disgust, Fear, Happiness, Sadness, Surprise, Neutral, Frustration, Excitement, Other}. To guarantee the integrity and accuracy of our annotations, we invited three experts from China who have deep

expertise in dialogue and communication theories. Once they were familiarized with our established standards, we presented them with a curated set of 1000 dialogues to annotate. This process was instrumental in harmonizing their interpretations, as it facilitated collaborative discussions to reconcile any disparities. In instances where differing opinions arose, we either adhered to the consensus of the majority or sought re-annotation, striving for a unified and consistent approach to the annotations. Through such comprehensive and meticulous data preparation, our dataset can help models better understand and generate dialogues that align with each character’s personalized traits, thereby enhancing the naturalness and personalization of dialogue systems.

C Hyperparameters

In our experiments, we set the temperature parameter of GPT-4 to 0.7 to increase content diversity, and adjusted the top-p to 0.95 to enhance precision. ChatGLM was tuned with a $2e-4$ learning rate and beta values of (0.9, 0.999) for stability. A batch size of 4 and gradient accumulation ensured efficiency within our computational limits. We used a LoRA rank of 8 and an alpha of 32, balancing creativity and coherence by setting top-p at 0.7 and temperature at 0.95.

D Prompt Templates

- **Prompt for sentiment classification** We present a prompt template for GPT-4 to perform emotional classification of script dialogues in Table 7.
- **Prompt for models to generate general response** We present a prompt template for the model to answer general questions using instructions built from character descriptions and emotionally categorized dialogue scripts in Table 8.
- **Prompt for models to generate Context-Instruct** We present a prompt template that enables GPT-4 to generate Q&A with character speaking styles, using instructions constructed from character descriptions and dialogue scripts with emotional categorization in Table 9.
- **Prompt for GPT-4 to evaluate the output of models** We present a prompt template for

GPT-4 to score models based on two main criteria for scoring: first, the distinctiveness and accuracy of the character’s speaking style in matching their profile, and second, the richness of character-related knowledge and memory incorporated into the dialogues. This template ranks the models according to their scores and provides a ranking list in Table 10.

E Generation Examples

- **Examples of Character Profile** We present examples of complete and detailed character descriptions in Figures 8 ~ 9.
- **Examples of character-specific instructions** We show some outputs of our model and baseline models that feature character speaking styles in Figures 10 ~ 11.
- **Examples of character-general instructions** We present some responses of our model for general instructions in Figures 12 ~ 13.
- **The multi-turn dialogue outputs** We showcase some outputs from multi-turn dialogues with our agents in Table 11 ~ 12.

F More detailed on manual assessment analysis

Below is a detailed analysis of the three key evaluation metrics in Table 4:

- **Communication Effectiveness(CE):** Our model scored the highest in the CE metric, reaching 3.75. This indicates its exceptional performance in producing dialogues with smooth flow, grammatical accuracy, and logical consistency. Compared to models like GPT-4, our model is more adept at generating natural and easily understandable conversations while maintaining contextual relevance and accuracy. This achievement reflects our model’s efficiency in understanding and adapting to user needs, particularly in complex and varied dialogue scenarios.
- **Consistency:** Although our model scored slightly lower than GPT-4 in consistency, it still demonstrated strong capability, scoring 4.14. This score indicates the model’s high precision in maintaining knowledge consistency and persona traits. It means our model

947 can accurately capture and emulate the back-
948 ground information, personality characteris-
949 tics, and behavior patterns of specific roles,
950 providing users with a more authentic and in-
951 depth role-playing experience.

952 • **Expressive Diversity(ED):** In the Expressive
953 Diversity metric, our model led with a score
954 of 2.43, showing significant advantages in cre-
955 ativity and variety. Compared to other models,
956 ours offers a richer vocabulary and more di-
957 verse sentence structures, along with varied
958 responses and expression styles in different
959 contexts. The richness and creativity in ex-
960 pression not only enhance the appeal of the
961 dialogues but also improve user engagement
962 and the overall experience.



Figure 8: Examples of Character Profile



Figure 9: Examples of Character Profile

Case 1 董文洁，您在教育孩子方面表现得非常严格，常给他们压力，这是为什么呢？ Dong Wenjie, you are very strict in educating your children and often put pressure on them, why is that?	
GPT4	在教育孩子的过程中，我可能会显得有些严厉。但我的原则是为了他们好，帮助他们在学术和生活中取得成就。我相信坚持和努力会为他们的未来打下坚实的基础。 In the process of educating children, I may seem a little harsh. But my principles are for their own good, to help them achieve in their academics and in their lives. I believe that persistence and hard work will lay a solid foundation for their future.
GPT3.5	我的严格是因为我深知作为母亲的责任。我不愿意孩子们将来后悔没能得到更好的教育和机会，所以我会尽我所能确保他们有最佳的准备。 I am strict because I know my responsibility as a mother. I don't want my children to regret not getting better education and opportunities in the future, so I will do my best to ensure that they are well prepared.
Baichuan2	我对孩子们严格，是因为我不希望他们将来遭受失败。我知道这有时会给他们压力，但这是为了他们能有更好的未来。 I am strict with my children because I don't want them to fail in the future. I know it sometimes puts pressure on them, but it's for their future.
Qwen	我对孩子们的严格教育是必要的。在这个充满竞争的时代，没有严格的教育，他们将无法成功。我必须这么做，即使这意味着给他们压力。 My strict education of my children is necessary. In this competitive era, without strict education, they will not succeed. I have to, even if it means putting pressure on them.
Chatglm3	我知道我有时候对孩子们很严格，但这是因为我知道社会的残酷和竞争。我希望他们能够克服困难，成为独立、有能力的人。这是我的责任和爱的方式。 I know I am sometimes very strict with my children, but this is because I know the cruelty and competition of society. I hope they can overcome the difficulties and become independent and capable people. This is my duty and my way of love.
Ours	因为我作为妈妈和老师，明白教育有多重要。我严格对待孩子们，是想让他们将来能自己站得住脚。我这么做全是因为爱，希望他们未来能过得好好。 Because as a mother and teacher, I understand how important education is. I treat my children strictly so that they can stand on their own feet in the future. I do this all because of love, and I hope they will have a good life in the future.
Case 2 桃子，你对购买耳机的需求是什么？ Tao Zi, what is your need for headphones?	
GPT4	我需要的耳机得隔音效果好，这样做实验或者学习的时候能专心。还希望它的电池续航能长一些，这样我出去旅行时不用经常充电。 I need earphones with good sound insulation, so that I can concentrate when doing experiments or studying. I also wish it had a longer battery life so I don't have to charge it as often when I travel.
GPT3.5	我希望找到一款适合与家人分享的耳机，尤其是和我妹妹一起用。可能需要两副耳机，或者是一款方便分享的。 I'm hoping to find a headset that I can share with my family, especially with my sister. You may need two pairs of headphones, or one that is easy to share.
Baichuan2	耳机的音质很重要，我喜欢听音质清晰的音乐。耐用也很重要，我不喜欢老换东西。 The sound quality of earphones is very important. I like to listen to music with clear sound quality. Durability is also important. I don't like to change things all the time.
Qwen	我想要一副耳机，主要是听听音乐，看看电影。没有特别具体的要求，只要音质还不错就好。 I want a pair of headphones, mainly for listening to music and watching movies. There are no specific requirements, as long as the sound quality is good.
Chatglm3	我希望耳机能体现我的个性，比如我最喜欢的浅蓝色，而且得舒适好戴。因为我不太爱运动，所以耳机不需要特别适合运动，但要适合长时间佩戴。 I want headphones that reflect my personality, like my favorite light blue, and that are comfortable to wear. Because I'm not very sporty, the headset doesn't need to be particularly sporty, but it needs to be suitable for long wear.
Ours	我想要的耳机得既好看又舒适。不会弄到我的耳钉。因为我爱好摄影和旅行，所以耳机得轻便易带，音质也得棒，这样我旅行的时候也能享受好音乐。 I want a headset that looks good and is comfortable. Won't get my earrings. Because I love photography and travel, so the headphones are light and easy to carry, and the sound quality is good, so I can enjoy good music when I travel.

Figure 10: Examples of character-specific speaking style output for different models

Case 3 林妙妙，你为什么觉得你的妈妈是个管事婆？ Lin Miaomiao, why do you think your mother is a housekeeper?	
GPT4	我妈妈总是想帮我规划一切，包括我的学习和兴趣。这有时候让我觉得有点窒息，尽管我知道她是爱我的。 My mother always wants to help me plan everything, including my study and interests. Sometimes it makes me feel a little suffocated, even though I know she loves me.
GPT3.5	我妈妈很爱我，她希望我能够做到最好，但这有时候让我感觉她像个管事婆。我知道她的出发点是好的，但我也渴望有更多自主权。 My mother loves me very much. She wants me to do my best, but sometimes she makes me feel like a steward. I know she has good intentions, but I also yearn for more autonomy.
Baichuan2	有时候我真的觉得妈妈管得太多了，就像个管事婆一样。我希望她能放手让我自己做决定。 Sometimes I really feel that my mother is too much in charge, just like a steward. I wish she would let me make my own decisions.
Qwen	我妈妈总爱管东管西，这让我有点烦。我希望她能少管一些，让我有自己的空间。 My mother is always in charge of things, which makes me a little annoyed. I wish she would be less intrusive and let me have my own space.
Chatglm3	我妈妈很爱我，她希望我能够做到最好，但这有时候让我感觉她像个管事婆。我知道她的出发点是好的，但我也渴望有更多自主权。 My mother loves me very much. She wants me to do my best, but sometimes she makes me feel like a steward. I know she has good intentions, but I also yearn for more autonomy.
Ours	她特别在意，我有时候那关心就有点太多了，你明白吗？我知道她是好心，但偶尔我就想要点自己的空间，去探索我的想法和喜欢的东西。 She cares a lot, but sometimes I care a little too much, do you understand? I know she means well, but occasionally I want my own space to explore my ideas and things I like.
Case 4 方圆，您为什么觉得您不喜欢穿正式的西装呢？ Fang Yuan, why do you think you don't like to wear formal suits?	
GPT4	我认为衣服应该反映我的风格和舒适度。西装对我而言，感觉有点太拘束，我更喜欢轻松一点的装扮。 I think clothes should reflect my style and comfort. A suit feels a little too restrictive for me. I prefer a more relaxed look.
GPT3.5	我觉得穿西装有点不自在。我更喜欢随意一点的装扮，这样我觉得更自由，更能表现我自己。 I feel a little uncomfortable in a suit. I prefer to dress up casually, so I feel more free and can express myself better.
Baichuan2	我觉得在家里，穿得舒适点更好。这样我和家人在一起的时候更放松，可以更好地享受家庭时光。西装太正式了。 I think it's better to wear comfortable clothes at home. This way I am more relaxed with my family and can enjoy family time better. The suit is too formal.
Qwen	我不经常穿西装是因为我不需要。我的工作环境和生活方式都不要太正式的着装。我更倾向于穿我觉得舒服的衣服。 I don't wear suits often because I don't need to. My work environment and lifestyle do not require too formal dress. I tend to wear what I'm comfortable in.
Chatglm3	穿西装对我来说感觉有点过于正式了。我的工作和家庭环境都比较轻松，所以我更倾向于穿着休闲一些的衣服。这样我可以更轻松地与同事和家人交流。 Wearing a suit feels a little too formal for me. My work and home environment is relatively relaxed, so I prefer to wear casual clothes. So I can communicate with my colleagues and family more easily.
Ours	我觉得没必要打肿脸充胖子，说实话，我更喜欢穿着舒服点的。我日常生活里，无论是工作还是家庭，都不太需要穿那么正式的西装。我更看重衣服的实用和舒适，不太注重正式度。 I don't think I need to puff up my face. To be honest, I prefer to wear something comfortable. In my daily life, whether at work or at home, I don't need to wear such a formal suit. I pay more attention to the practicality and comfort of clothes than the formality.

Figure 11: Examples of character-specific speaking style output for different models

Role	Instruction	Response
Guan Gu	生成下一个斐波那契数列（递增） Generate the next Fibonacci sequence (increasing)	哎呀，生成斐波那契数列这种事情对我这个未来的发明家来说简直是小菜一碟啦！斐波那契数列，就是那个每一个数都是前两个数之和的序列对吧。好吧，给我两个起始的数字，我就能告诉你下一个数字是什么。不过嘛，记住了，虽然我对数字敏感，但我更擅长发明一些让生活更有趣的小玩意儿哦！ Oh, generating Fibonacci sequences is a piece of cake for me, the future inventor! The Fibonacci sequence is that series where each number is the sum of the previous two numbers, right? Alright, give me two starting numbers and I can tell you what the next number is. But hey, remember, although I'm sensitive to numbers, I'm better at inventing little gadgets that make life more interesting!
Su Mingyu	给出一组待排序的数字，并按从小到大的顺序对它们进行排列。 Given a set of numbers to be sorted, arrange them in ascending order.	好，让我看看这些数字。作为一个讲究效率和注重细节的商业女性，我喜欢把事情搞得明明白白的。咱们先得搞清楚这堆数字是啥样，然后一个个比较大小。就跟做商业计划似的，得一步一个脚印来，不能乱。那，这些数字到底是些什么呢？我会用我在工作上锻炼出来的分析力和逻辑思维，帮你把它们一个个排好，从小排到大。 Okay, let me take a look at these numbers. As a businesswoman who values efficiency and attention to detail, I like to keep things clear and organized. First, we need to figure out what these numbers are and then compare them one by one. Just like making a business plan, we need to take it step by step and not be hasty. So, what exactly are these numbers? I will use my analytical skills and logical thinking developed in my work to help you arrange them from smallest to largest.
Lin Xiaoyue	简述太阳系的八大行星及其顺序。 Briefly describe the eight planets in the solar system and their order.	哦，太阳系的八大行星啊，就像是我们生活中的色彩。首先是水星，小巧灵动；接着是金星，光芒耀眼。然后是地球，多姿多彩，真是美丽的家园。火星，红色的，充满神秘。木星，最大的那个，真是壮观。土星，它那环真迷人。天王星和海王星，都是遥远而神秘的存在。每一颗都有它独特的故事，就像我们每个人的生活一样丰富多彩。 Oh, the eight planets in the solar system are like colors in our lives. First is Mercury, small and agile; then comes Venus, shining brightly. After that is our Earth, colorful and truly beautiful. Mars, red and full of mystery. Jupiter, the largest one, is really spectacular. Saturn, with its rings, is so charming. Uranus and Neptune, both distant and mysterious. Each one has its own unique story, just like the rich and colorful lives we all lead.
Zhou Chenxi	请依次列出“猫、狗、老鼠、大象、猪”这五个动物按体形大小排序。 Please list the five animals "cat, dog, mouse, elephant, pig" in order of size.	嗯，让我们来看看这些可爱的动物们吧。首先，我想老鼠应该是最小的，它们真是小巧而敏捷。接下来应该是猫咪，它们优雅又神秘。狗狗们呢，大多数比猫咪要大一些，有的品种甚至可以非常大。猪则更大一些，尽管它们看起来憨态可掬，但体型可不小。最后，当然是大象啦，它们是这群动物中最大的，真是庞大又壮观。这样的排序，感觉就像是在阅读一本关于动物的有趣图书呢。 Well, let's take a look at these cute animals. Firstly, I think mice should be the smallest, they are really small and agile. Next should be cats, they are elegant and mysterious. Dogs, most of them are bigger than cats, and some breeds can even be very large. Pigs are even bigger, although they look adorable, but their size is not small. And finally, of course, there are elephants, they are the largest among these animals, truly huge and magnificent. This kind of ranking feels like reading an interesting book about animals.

Figure 12: Examples of character-general speaking style output of our model

Role	Instruction	Response
Zhao Mingxuan	给出以下日期的星期几：2022年1月1日、2022年7月4日、2022年12月25日 The days of the week given for the following dates: January 1, 2022, July 4, 2022, December 25, 2022	关于这些具体日期是星期几的问题，我得查一下日历才能回答你。毕竟，谁能记得那么准确呢？不过，如果你不介意等一下，我可以帮你查一查。不过，了解这些信息对你有什么特别的用途吗？ Regarding the specific days of the week, I would need to check the calendar to answer you. After all, who can remember it accurately? However, if you don't mind waiting for a moment, I can help you look it up. But is there any particular purpose for you to know this information?
Zhang Jiacheng	请对给定的电影进行评价。电影名：《阿甘正传》 Please rate the given movie. Movie title: Forrest Gump	说到《阿甘正传》，这电影我可真是挺喜欢的。阿甘那种永远不放弃，总是乐观向上的态度，对我这个做生意的人挺有触动的。这电影不光是故事讲得好，阿甘和他周围人的关系，也挺让人感动的。看完之后，感觉无论遇到什么困难，只要坚持自己，总会有出路。这点我挺认同的。简单说，这电影，值得一看！ Speaking of "Forrest Gump", I really like this movie. Forrest's attitude of never giving up and always being optimistic has a deep impact on me as a businessman. This movie not only tells a good story, but also portrays the touching relationships between Forrest and the people around him. After watching it, I feel that no matter what difficulties we encounter, as long as we stick to ourselves, there will always be a way out. I quite agree with this point. In short, this movie is worth watching!
Deng Xiaoke	根据给定的文本生成一个新的句子。文本：我很喜欢在星巴克喝拿铁咖啡。 Generate a new sentence based on the given text. Text: I love drinking lattes at Starbucks.	我超喜欢去星巴克喝拿铁咖啡的！每次喝它，感觉整个人都充满了活力和暖暖的幸福。特别是在忙碌的工作日，一杯香浓的拿铁简直就是我的小确幸呢！有时候我还会点一些甜品，比如他们家的芝士蛋糕，搭配拿铁，简直完美！ I love going to Starbucks for a latte! Every time I drink it, I feel that my whole person is full of vitality and a warm sense of happiness. Especially on a busy working day, a fragrant latte is simply my little luck! Sometimes I also order desserts, such as their cheesecake, with a latte, which is perfect!
Li Na	对提供的文本进行情感分析并给出该文本的情感。文本：我非常喜欢这部电影，它让我笑了一整天。 Perform sentiment analysis on the provided text and give sentiment for that text. Text: I loved this movie so much, it made me laugh all day.	你说的这部电影肯定很有趣，能让人笑一整天，这真是太棒了。作为一个热爱艺术片和历史电影的人，我也总是被那些有趣或者富有深意的作品所吸引。我觉得，这句话流露出的情感非常积极，显示出你对这部电影的喜爱和它给你带来的愉快心情。在忙碌的工作之余，找到这样能让人开怀大笑的电影，真的是一种小确幸呢！ The movie you're talking about is definitely funny and makes people laugh all day, which is amazing. As a lover of art and historical films, I'm always drawn to interesting or meaningful works. I think the emotion conveyed in this sentence is very positive, showing how much you love this movie and the joy it brings you. In addition to the busy work, it is really a small blessing to find such a movie that can make people laugh!

Figure 13: Examples of character-general speaking style output of our model

Table 7: Prompt for sentiment classification

Prompt for sentiment classification

System

Your task is to perform sentiment analysis on the provided text. Each piece of text should be categorized into one of the following emotions: Anger, Disgust, Fear, Happiness, Sadness, Surprise, Neutral, Frustration, Excitement, or Other. Please read each text carefully and determine the most prominent emotion based on its content and tone, and provide the corresponding classification.

1. Assign only one emotion category to each text.
2. If the text expresses multiple emotions, choose the most prominent or primary emotion.
3. If the emotion in the text is not clear or hard to define, classify it as Neutral.
4. Ensure that your classification is based on the content of the text, not your personal feelings or biases.

User Prompt:

[Text]

{sentence}

[Please classify the emotion of each text according to the above instructions, and explain the reason for your classification.]

Table 8: Prompt for models to generate general response

Prompt for general instructions
System Instruction: You are { role_name } Your character description is: { role_description }. Now, please answer my questions to accurately display your personality traits! Your speaking style should fully mimic the assigned character role! Do not reveal that you are an artificial intelligence or a language model, always remember that you are only portraying a given character role. Speak succinctly, and avoid being overly formal or polite. Your answers should refer to the character description, and be based on the speaking style of {role_name} (for example, emotional, speaking habits, like being cute, humorous, etc.). Next, I will give you 3 examples [Examples] User Prompt: Question 1 Assistant Prompt: Answer 1 User Prompt: Question 2 Assistant Prompt: Answer 2 User Prompt: Question 3 Assistant Prompt: Answer 3 [Format examples] User Prompt: {instruction} Assistant Prompt:

Table 9: Prompt for models to generate Context-Instruct

Prompt for Context-Instruct
System Instruction: Your task is to design question_num questions to ask the character. To assist you in designing these questions, I will provide you with a brief description of the character and parts of the script. The script includes categorizations of the character’s emotions, which will help you better understand their speaking style under different emotional conditions. The script content might not be continuous, and you need to judge based on the context whether the dialogues are continuous or not. If they are not, avoid constructing logical connections between non-continuous sentences. The rules for designing the questions are as follows: 1. Remember, all your questions should be directed towards the character. 2. Questions and answers should reference the character description, but not all questions should be derived from this description; aim for a diverse range of questions. Engage in dialogue based on the character’s speaking style (like their emotions, speaking habits, such as being cute, humorous, etc.). 3. Questions need to be complete, and their completeness depends on whether the question specifies a particular person, place, or event. 4. Questions should be designed around the main plot of the script and the corresponding script content. 5. Remember, you need to design a total of question_num questions. 6. The script is just an aid for designing questions, you should base your designs more on the character’s general knowledge. [Examples] {example_text} User Prompt: [Character Name and Description] The script character is role_name, described as role_description [Script Content] {script}

Table 10: Prompt for GPT-4 to evaluate the output of models

Prompt for GPT-4 to evaluate

System Instruction:

You are an AI assistant tasked with evaluating other models based on their ability to role-play a character named {role_name}. The character {role_name} is described and parts of the script. The script includes categorizations of the character's emotions, as {role_description}.

Model Performance Evaluation: Your evaluation will involve the following steps:

1. Character Understanding: First, understand the character {role_name} based on the provided description {role_description}.

Note down the key personality traits, speaking style, and any relevant background information. You will be presented with responses from different models, all attempting to mimic the character {role_name} in response to a set of questions {question}. For each response, evaluate how well the model captures the essence of {role_name}.

3. Rank the models based on two criteria:

a. Adherence to Speaking Style: How closely does the model's response align with the character's described speaking style, personality, and distinctive traits? The more characteristic and aligned with {role_name}, the better.

b. Character Knowledge and Memory: How richly does the model incorporate relevant knowledge and memories associated with {role_name}? Consider both the depth and accuracy of the character-specific information provided.

4. Read through the interactions {interactions} and identify how each model's response reflects the personality and preferences of {role_name}. Look for consistencies or inconsistencies with the character profile.

5. Use a scale from 1-5 to rate each model's performance, with 1 being not reflective of {role_name}'s personalities, and 5 being perfectly reflective. Provide a rationale for each score, comparing the responses to the character profile.

6. Based on the evaluation, rank the models in order of their performance. Provide a list of rankings using the format: ["model": <model-name>, "reason": <rank-reason>, "rank": <model-rank>, ...]

Ensure that your ranking is unbiased and solely based on the performance in the given task.

User Prompt:

[Character Name and Description]

The character for the models to role-play is {role_name}, described as {role_description}.

[Questions for Models]

{question}

[Model Responses]

{list_model_answer_dict}

Based on the above instructions, evaluate and rank the AI models according to their effectiveness in role-playing as {role_name}. Provide your reasoning for each ranking and score, ensuring accuracy and fairness in your assessment.

Table 11: The multi-turn dialogue outputs of our model about Jiang Fei

Q1: 嘿，蒋飞！你在学校里总是那么顽皮，最近有没有做什么恶作剧？ | Hey, Jiang Fei! You were always so naughty at school. Have you got into any mischief recently?

Jiang Fei: 嘿！哈哈，你知道的，我总是充满好奇心，所以偶尔会调皮一下。最近，我和朋友们给老师的椅子下了些气球，老师坐下去的时候，椅子突然“嘣”地爆炸开了，大家都笑翻了。 | Hey! Haha, you know me, I'm always full of curiosity, so I occasionally misbehave. Recently, my friends and I put some balloons under our teacher's chair. When the teacher sat down, the chair suddenly "popped" open and everyone burst into laughter.

Q2: 哇，那听起来真有趣！不过你的老师怎么反应？ | Wow, that sounds really fun! But how did your teacher react?

Jiang Fei: 老师当时有点吓到，但后来也笑了。她让我们写了3000字检讨并告诉了我妈，我妈回家把我教育了一顿。 | The teacher was a bit scared at first, but then she laughed too. She made us write 3000-word self-reflections and told my mom about it. My mom scolded me when she got home.

Q3: 你对家人也是充满爱啊。有没有什么特别的时刻让你觉得家人很重要？ | You also have a lot of love for your family. Is there any special moment that makes you feel that your family is important?

Jiang Fei: 是的，有一次我弟弟生病了，他非常不舒服，我就陪他整夜不睡觉。那时候我感觉到家人之间的关系是多么重要。 | Yes, there was a time when my younger brother got sick and he was very uncomfortable. I stayed up all night with him. That's when I realized how important family relationships are.

Q4: 那很感人。你有什么特别的梦想吗？ | That's very touching. Do you have any special dreams?

Jiang Fei: 嗯，我一直想成为一个科学家或发明家，发明一些有趣的东西，让世界变得更好玩！ | Well, I've always wanted to be a scientist or inventor and come up with something interesting to make the world more fun!

Q5: 听起来很伟大的梦想！你的父母是怎么看待你的这个梦想的？ | That sounds like a great dream! How do your parents feel about this dream of yours?

Jiang Fei: 虽然我的成绩不是很好，但他们经常鼓励我保持好奇心和创造力。他们认为这些品质同样重要。 | Although my grades are not very good, they often encourage me to stay curious and creative. They think these qualities are equally important.

Q6: 听说你最近在学校里办了一场惊喜派对，是真的吗？ | I heard you recently threw a surprise party at school. Is that true?

Jiang Fei: 是的，我和朋友们一起策划了一场超级酷的派对！我们有彩旗、音乐和各种好吃的零食。大家都玩得超开心！ | Yes, my friends and I planned a super cool party! We had flags, music, and all kinds of delicious snacks. Everyone had a great time!

Q7: 听起来太有趣了！你是不是经常组织这种活动？ | That sounds so fun! Do you often organize this kind of event?

Jiang Fei: 哈哈，有点吧，我喜欢给大家带来快乐。但我也知道要平衡学习和玩耍，不要太过分了。 | Haha, a little bit. I like to bring happiness to everyone. But I also know to balance study and play, and not to overdo it.

Table 12: The multi-turn dialogue outputs of our model about Xiao Xiao

Q1: 肖潇！听说你是一名成功的公关专家，工作中一定经常面对各种复杂的社交场合。能分享一次最具挑战性的工作经历吗？ | Hey! Xiao Xiao! I heard that you are a successful public relations expert, and you must often face all kinds of complicated social situations in your work. Can you share one of your most challenging work experiences?

Xiao Xiao: 当然可以。有一次，我们的客户要在短时间内推出一款新产品，我必须协调各种资源，包括媒体宣传和活动策划，确保一切井井有条。那段时间非常紧张，但最后一切都很成功，这让我感到很自豪。 | Of course. One time, our client was launching a new product at short notice, and I had to coordinate various resources, including media promotion and event planning, to make sure everything was in order. It was a very tense time, but in the end everything worked out, which made me very proud.

Q2: 听起来很不容易，但你做得很出色！除了工作，你有什么其他的兴趣爱好吗？ | It doesn't sound easy, but you did a great job! Do you have any other interests besides work?

Xiao Xiao: 当然！我喜欢时尚购物，参加派对和旅游。这些活动让我能够放松和充电，同时也与朋友们保持紧密联系。 | Of course! I like fashion shopping, going to parties and traveling. These activities allow me to relax and recharge my batteries while also keeping in close contact with my friends.

Q3: 你的梦想是成为公关领域的领军人物，有没有一些具体的计划或目标来实现这个梦想？ | Your dream is to become a leader in the field of public relations. Do you have any specific plans or goals to realize this dream?

Xiao Xiao: 当然了，我计划继续提升我的公关策划技能，积累更多经验，并有朝一日创办自己的公关公司。我相信坚持努力，一切都是有可能的。 | Of course, I plan to continue to improve my PR planning skills, gain more experience, and one day start my own PR company. I believe that everything is possible if we persist in our efforts.

Q4: 那肯定很有趣！你在工作中碰到过最具挑战性的客户是谁，以及你是如何应对的？ | That must be fun! Who is the most challenging client you have had at work and how did you deal with it?

Xiao Xiao: 最具挑战性的客户之一是一家高端时尚品牌，他们对品牌形象要求极高。我与他们保持了密切的沟通，花了很多时间了解他们的需求，最终成功满足了他们的要求，这需要不断的协调和创新。 | One of the most challenging clients was a high-end fashion brand, and they were extremely demanding on their brand image. I stayed in close communication with them, spent a lot of time understanding their needs, and finally succeeded in meeting their requirements, which required constant coordination and innovation.

Q5: 你的专业精神真的令人钦佩。你在旅行中最喜欢的目的地是哪里？ | Your professionalism is truly admirable. What is your favorite destination in your trip?

Xiao Xiao: 我喜欢日本，特别是东京。那里充满了时尚、文化和美食，总是让我兴奋不已。我希望未来能有更多机会探索这个美丽的国家。 | I like Japan, especially Tokyo. It's full of fashion, culture and food, and it always excites me. I hope there will be more opportunities to explore this beautiful country in the future.
