



CIKM SEOUL 2025
November 10 - 14



Multi-Personality Generation of LLMs at Decoding-time

Rongxin Chen^{1,2}, Yunfan Li^{1,2}, Yige Yuan^{1,2}, Bingbing Xu¹, Huawei Shen^{1,2}

¹Institute of Computing Technology, Chinese Academy of Sciences

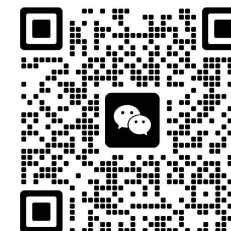
²University of Chinese Academy of Sciences



Paper



WeChat



Presenter: Rongxin Chen

Email: chenrongxin24s@ict.ac.cn



CIKM SEOUL
2025
November 10 - 14

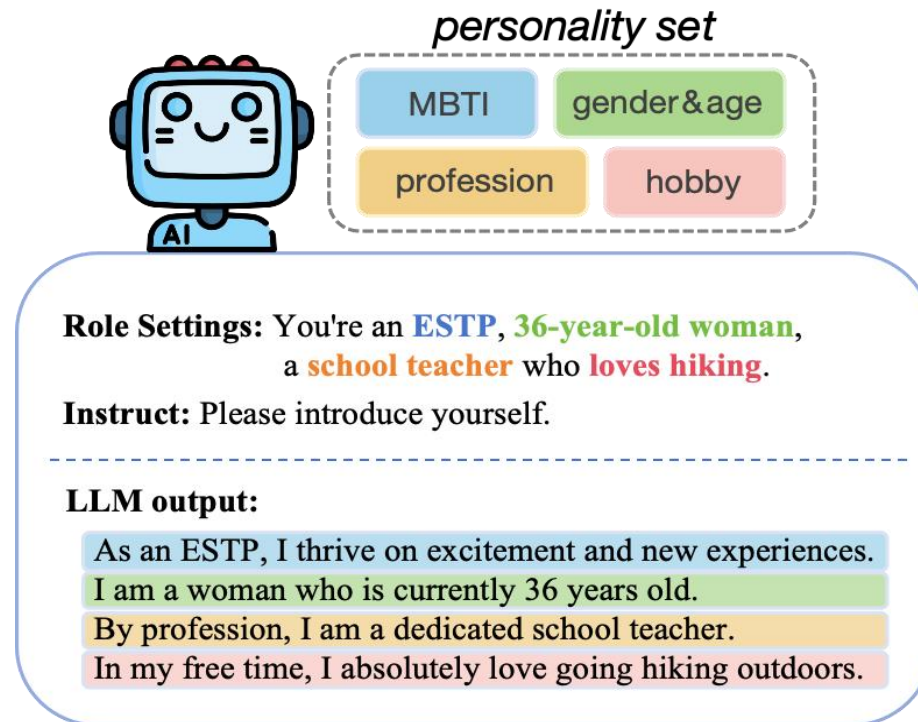


- **Motivation**
- Method
- Experiments

MOTIVATION

Multi-Personality Generation

- Multi-personality generation requires generated text to simultaneously embody **multiple, potentially interacting** personalization attributes.



MOTIVATION

Existing Work

- Retraining-based methods
 - Use multi-objective optimization (e.g., MORLHF) to "integrate" multiple preferences into a single model's weights.
 - Costly, scale poorly, and struggle to adapt to new preferences.

(Base LLM) + (Prefs A, B, C)
——> [Expensive Retraining] ——> (Model_ABC)

MOTIVATION

Existing Work

- Decoding-time methods
 - Guide the decoding process using external models (e.g., reward models) or by mixing logits from multiple models.
 - Multi-dimensional reward models are scarce and hard to acquire.
 - Simple model combination (e.g., logits aggregation) mechanisms are heuristic and lack robustness.

(Model A) + (Model B) + (Model C)
————> [Logits Mixer] ———> (Output Token)

MOTIVATION

Main Idea

- **Observation:** When we align a model (e.g., for "Hobby") from a base model using DPO, all the preference information is encoded.
- **Where?** In the Probability Density Ratio: $r(y|x) = \frac{\pi_{\theta}(y|x)}{\pi_{\text{ref}}(y|x)}$
- This ratio is a "free lunch" —it quantifies the exact preference of a single attribute, and we already have it.
- **Main Idea:** *Can we skip retraining and complex reward models, and instead **directly aggregate these "free" density ratios** from multiple single-dimensional models at decoding-time?*



CIKM ^{SEOUL} 2025
November 10 - 14



- Motivation
- **Method**
- Experiments

PROBLEM DEFINITION

Multi-Personality Generation

- **Goal:** We want to sample from a target distribution that aggregates N preferences.
- **Given:** A reference model: $\pi_{ref}(y|x)$
N single-preference models: $\{\pi_{d_i}(y|x)\}_{i=1}^N$
- **Target Policy:** Define the optimal policy π_{MPG} as one that optimizes a combined reward.

$$\pi_{MPG}(y|x) = \frac{\pi_{ref}(y|x)}{Z(x)} (\nabla f)^{(-1)} \left(\frac{1}{\beta} \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{ref}(y|x)} \right) \right)$$

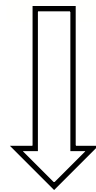
METHOD: MPG Framework

The Bridge: MPG & Rejection Sampling

Target Policy

$$\pi_{\text{MPG}}(y|x) = \frac{\pi_{\text{ref}}(y|x)}{Z(x)} (\nabla f)^{(-1)} \left(\frac{1}{\beta} \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right) \right)$$

ratio(target and proposal)
in **Rejection Sampling**



$$\boxed{\frac{\pi_{\text{MPG}}(y|x)}{\pi_{\text{ref}}(y|x)}} = \frac{1}{Z(x)} (\nabla f)^{(-1)} \left(\frac{1}{\beta} \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right) \right)$$

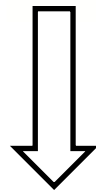
METHOD: MPG Framework

The Bridge: MPG & Rejection Sampling

Target Policy

$$\pi_{\text{MPG}}(y|x) = \frac{\pi_{\text{ref}}(y|x)}{Z(x)} (\nabla f)^{(-1)} \left(\frac{1}{\beta} \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right) \right)$$

ratio(target and proposal)
in **Rejection Sampling**



$$\boxed{\frac{\pi_{\text{MPG}}(y|x)}{\pi_{\text{ref}}(y|x)}} = \frac{1}{Z(x)} (\nabla f)^{(-1)} \left(\frac{1}{\beta} \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right) \right)$$

Acceptance Probability

$$A(y|x) = \boxed{\frac{\pi_{\text{MPG}}(y|x)}{M \cdot \pi_{\text{ref}}(y|x)}}$$

METHOD: MPG Framework

The Bridge: MPG & Rejection Sampling

Target Policy

$$\pi_{\text{MPG}}(y|x) = \frac{\pi_{\text{ref}}(y|x)}{Z(x)} (\nabla f)^{(-1)} \left(\frac{1}{\beta} \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right) \right)$$

ratio(target and proposal)
in **Rejection Sampling**

$$\boxed{\frac{\pi_{\text{MPG}}(y|x)}{\pi_{\text{ref}}(y|x)}} = \frac{1}{Z(x)} (\nabla f)^{(-1)} \left(\frac{1}{\beta} \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right) \right)$$

$$\frac{\pi_{\text{MPG}}(y|x)}{\pi_{\text{ref}}(y|x)} \propto \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right)$$

Acceptance Probability

$$A(y|x) = \boxed{\frac{\pi_{\text{MPG}}(y|x)}{M \cdot \pi_{\text{ref}}(y|x)}}$$

METHOD: MPG Framework

The Bridge: MPG & Rejection Sampling

Target Policy

$$\pi_{\text{MPG}}(y|x) = \frac{\pi_{\text{ref}}(y|x)}{Z(x)} (\nabla f)^{(-1)} \left(\frac{1}{\beta} \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right) \right)$$

ratio(target and proposal)
in Rejection Sampling

$$\boxed{\frac{\pi_{\text{MPG}}(y|x)}{\pi_{\text{ref}}(y|x)}} = \frac{1}{Z(x)} (\nabla f)^{(-1)} \left(\frac{1}{\beta} \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right) \right)$$

Acceptance Probability

$$A(y|x) = \boxed{\frac{\pi_{\text{MPG}}(y|x)}{M \cdot \pi_{\text{ref}}(y|x)}}$$

$$\frac{\pi_{\text{MPG}}(y|x)}{\pi_{\text{ref}}(y|x)} \propto \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right) \xrightarrow[\nabla f(x) = \log x + 1]{f(x) = x \log x} \frac{\pi_{\text{MPG}}(y|x)}{\pi_{\text{ref}}(y|x)} \propto \sum_{i=1}^N \alpha_i (\log r_i(y|x) + 1) \propto \sum_{i=1}^N \alpha_i r_i(y|x)$$

METHOD: MPG Framework

The Bridge: MPG & Rejection Sampling

Target Policy

$$\pi_{\text{MPG}}(y|x) = \frac{\pi_{\text{ref}}(y|x)}{Z(x)} (\nabla f)^{(-1)} \left(\frac{1}{\beta} \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right) \right)$$

ratio(target and proposal)
in Rejection Sampling

$$\boxed{\frac{\pi_{\text{MPG}}(y|x)}{\pi_{\text{ref}}(y|x)}} = \frac{1}{Z(x)} (\nabla f)^{(-1)} \left(\frac{1}{\beta} \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right) \right)$$

$$\frac{\pi_{\text{MPG}}(y|x)}{\pi_{\text{ref}}(y|x)} \propto \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right)$$

$$\begin{array}{c} f(x) = x \log x \\ \hline \nabla f(x) = \log x + 1 \end{array} \rightarrow$$

$$\frac{\pi_{\text{MPG}}(y|x)}{\pi_{\text{ref}}(y|x)} \propto \sum_{i=1}^N \alpha_i (\log r_i(y|x) + 1) \propto \sum_{i=1}^N \alpha_i r_i(y|x)$$

Acceptance Probability

$$A(y|x) = \boxed{\frac{\pi_{\text{MPG}}(y|x)}{M \cdot \pi_{\text{ref}}(y|x)}}$$

$$\boxed{A(y|x) = \frac{\pi_{\text{MPG}}(y|x)}{M \cdot \pi_{\text{ref}}(y|x)} \propto \frac{\sum_{i=1}^N \alpha_i r_i(y|x)}{M}}$$

METHOD: MPG Framework

The Bridge: MPG & Rejection Sampling

Target Policy

$$\pi_{\text{MPG}}(y|x) = \frac{\pi_{\text{ref}}(y|x)}{Z(x)} (\nabla f)^{(-1)} \left(\frac{1}{\beta} \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right) \right)$$

ratio(target and proposal)
in Rejection Sampling

$$\boxed{\frac{\pi_{\text{MPG}}(y|x)}{\pi_{\text{ref}}(y|x)}} = \frac{1}{Z(x)} (\nabla f)^{(-1)} \left(\frac{1}{\beta} \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right) \right)$$

Acceptance Probability

$$A(y|x) = \boxed{\frac{\pi_{\text{MPG}}(y|x)}{M \cdot \pi_{\text{ref}}(y|x)}}$$

$$A(y|x) = \frac{\pi_{\text{MPG}}(y|x)}{M \cdot \pi_{\text{ref}}(y|x)} \propto \frac{\sum_{i=1}^N \alpha_i r_i(y|x)}{M}$$

The acceptance probability can be regarded as
just a weighted sum of ratios!

$$\frac{\pi_{\text{MPG}}(y|x)}{\pi_{\text{ref}}(y|x)} = \frac{1}{Z(x)} (\nabla f)^{(-1)} \left(\frac{1}{\beta} \sum_{i=1}^N \alpha_i \nabla f \left(\frac{\pi_{d_i}(y|x)}{\pi_{\text{ref}}(y|x)} \right) \right) \quad \alpha_i r_i(y|x)$$

SCR: The Efficient Algorithm

Speculative Chunk-level based Rejection sampling

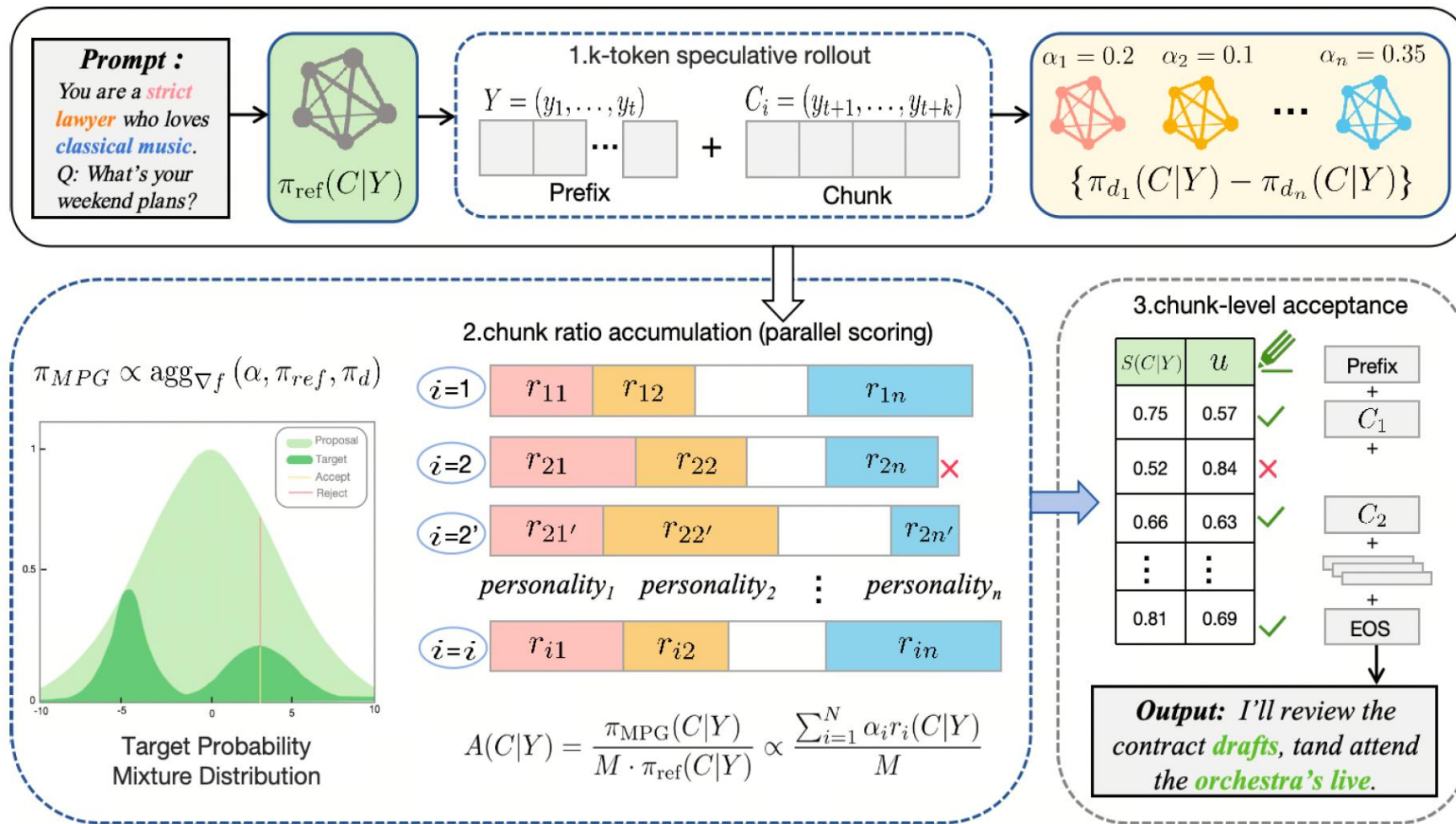
- **How it works:** full sequences $\implies$ **parallel** chunks (k-tokens)
proposal and validation

1. Propose: ref_model generates a k-token chunk C.
2. Validate: Get the score for that chunk $\text{Score}(C \mid Y) = \sum_{i=1}^N \alpha_i r_i(C \mid Y)$ by querying all single-preference models.
3. Accept: accept/reject the chunk using a dynamic threshold M. And **Online Update of M** via a three-phase sliding window:

Warm-up $\rightarrow$ Estimation $\rightarrow$ Stabilization

SCR: The Efficient Algorithm

Speculative Chunk-level based Rejection sampling





CIKM ^{SEOUL} 2025
November 10 - 14



- Motivation
- Method
- **Experiments**

EXPERIMENTS

Experimental Settings

2 tasks: MBTI Personality Simulation and Role-Playing

- **Datasets**

- MBTI-QA: instruction-following QA
- MBTI-MCQA: multiple-choice answering
- MBTI-16P: 16 Personalities psychometric instrument

- **Evaluation**

- **LLM-as-a-Judge**

- **Metrics:** Style (Sty), Thought (Tho), Behavior (Beh), Naturalness (Nat)

- **Datasets**

- Adapt the methodology from ALOE¹, focusing on user profile and personality.

- **Evaluation**

- **LLM-as-a-Judge**

- **Metrics:** Profile Relevance (PR), Persona Match (PM), Humanlikeness (HI)

- **Reference-based**

- **Metrics:** BLEU, ROUGE-1 F1, BERTScore, Perplexity

¹<https://github.com/ShujinWu-0814/ALOE>

EXPERIMENTS

Experimental Settings

2 tasks: MBTI Personality Simulation and Role-Playing

- **Baselines**

- **Base:** Llama-3-8B-Instruct.
- **DPO:** Direct Preference Optimization applied to the base model.
- **Preference Prompting (PP):** implements personality conditioning through explicit attribute descriptions in prompt engineering, applied across base model, single-preference model, and specialized model configurations.
- **DPO Soups (a variant of Personalized Soups):** acquires the LoRA parameters of independently - trained single - preference DPO models. Following weighted averaging, these parameters are applied to π_{ref} to generate text with multi-personality.
- **MOD:** performs a weighted linear combination of the output logits from single - attribute DPO models during decoding and samples based on the combined logits.

EXPERIMENTS

Main results

Method	QA					MCQA					16P					Overall
	Sty	Tho	Beh	Nat	Avg	Sty	Tho	Beh	Nat	Avg	Sty	Tho	Beh	Nat	Avg	
Evaluated by GPT-4o																
Base	3.282	3.722	3.634	3.166	3.451	2.688	3.069	3.014	2.473	2.811	2.851	3.115	3.092	2.745	2.989	3.084
Preference Prompting	3.348	3.796	3.657	3.341	3.535	2.770	<u>3.168</u>	3.131	2.538	2.902	2.888	3.273	3.287	2.854	3.076	3.171
DPO(single)	3.559	3.804	3.885	4.003	3.813	1.911	2.135	2.053	2.079	2.044	3.149	3.543	<u>3.442</u>	3.809	3.486	3.114
DPO Soups	3.458	3.937	3.850	3.494	3.685	2.751	3.161	<u>3.196</u>	2.557	<u>2.916</u>	2.936	3.203	3.345	2.892	3.094	3.232
MOD	3.517	<u>4.049</u>	<u>3.911</u>	<u>4.066</u>	<u>3.886</u>	1.727	1.846	1.758	1.946	1.819	<u>3.272</u>	<u>3.692</u>	3.431	<u>3.864</u>	<u>3.565</u>	3.090
SCR(ref-Base)	<u>3.583</u>	4.030	3.878	<u>3.677</u>	3.792	2.832	3.305	3.402	<u>2.692</u>	3.058	2.933	3.317	3.395	3.008	3.163	<u>3.338</u>
SCR(ref-DPO single)	3.927	4.322	4.444	4.271	4.241	<u>2.789</u>	3.026	2.866	2.803	2.846	3.481	3.900	3.786	4.000	3.792	3.626
Evaluated by DeepSeek-R1																
Base	3.375	4.025	3.978	3.404	3.696	2.617	3.523	3.602	2.598	3.085	2.952	3.546	3.759	3.035	3.323	3.368
Preference Prompting	3.438	4.093	4.139	3.508	3.794	2.785	3.710	3.777	2.697	3.242	3.058	3.706	3.831	3.172	3.442	3.493
DPO(single)	3.792	3.989	4.113	4.263	4.039	2.016	2.160	2.185	2.115	2.119	3.570	3.771	3.921	4.346	3.902	3.353
DPO Soups	3.569	<u>4.296</u>	4.335	3.768	3.992	2.821	<u>3.810</u>	3.968	2.848	<u>3.362</u>	3.117	<u>3.911</u>	<u>4.119</u>	3.217	3.591	3.648
MOD	<u>3.825</u>	<u>4.077</u>	4.207	<u>4.399</u>	<u>4.127</u>	1.746	1.798	1.830	1.878	1.813	<u>3.691</u>	3.892	3.941	<u>4.355</u>	<u>3.970</u>	3.303
SCR(ref-Base)	<u>3.656</u>	4.286	<u>4.337</u>	3.833	4.057	<u>2.876</u>	3.824	<u>3.871</u>	<u>2.912</u>	3.373	3.147	<u>3.911</u>	3.992	3.314	3.591	<u>3.674</u>
SCR(ref-DPO single)	4.153	4.401	4.713	4.584	4.463	2.920	3.501	3.244	3.084	3.137	3.900	4.153	4.147	4.422	4.156	3.919

Table: Comparison of baseline methods and SCR on MBTI task. SCR(ref-Base) and SCR(ref-DPO single) refer to the use of Base and DPO single, respectively, as the reference model to perform the SCR.

EXPERIMENTS

Main results

Method	Evaluated by GPT-4o				Evaluated by DeepSeek-R1				Reference-based Evaluation			
	PR	RM	HI	Avg	PR	RM	HI	Avg	BLEU	ROGUE-1	BERTScore	PPL
Base	3.580	3.770	3.471	3.607	3.418	3.778	4.098	3.765	0.010	0.088	0.762	43.984
Preference Prompting	3.720	3.853	3.832	3.802	3.582	3.821	4.283	3.895	0.024	0.127	0.823	<u>41.860</u>
DPO(single)	3.823	4.240	4.137	4.066	3.744	4.106	4.563	4.137	0.058	0.167	0.868	48.262
DPO Soups	3.818	3.998	3.871	3.896	3.692	3.857	4.379	3.976	0.022	0.130	0.823	43.594
MOD	<u>4.020</u>	4.298	<u>4.170</u>	<u>4.166</u>	<u>3.970</u>	4.187	<u>4.586</u>	4.241	<u>0.082</u>	<u>0.187</u>	<u>0.870</u>	47.188
SCR(ref-Base)	3.874	3.997	3.861	3.911	3.722	3.833	4.414	3.990	0.034	0.133	0.829	36.694
SCR(ref-DPO single)	4.030	<u>4.290</u>	4.192	4.167	3.976	<u>4.166</u>	4.596	<u>4.230</u>	0.085	0.188	0.897	43.335

Table: Comparison of baseline methods and SCR on Role-Playing task.

Experiments on 2 tasks indicate the effectiveness of MPG, SCR outperforms baselines, showing improvements up to 16%–18%.

EXPERIMENTS

Analysis 1: Efficiency (SCR is Fast)

Method	Throughput $\uparrow$ (tok/s)	Latency $\downarrow$ (s/seq)	FwdPass $\downarrow$ (per tok)	Reject Rate $\downarrow$ (%)
Base	120.0	1.07	1.0	—
DPO Soups	118.0	1.08	~ 1.0	—
MOD	60.0	2.13	2.0	—
Seq-RS	11.0	11.64	10.0	60.0
Token-RS	30.0	4.27	4.62	25.0
SCR	97.0	1.20	1.4	30.0

Table: Efficiency comparison of different methods for Role-Playing task ($N=2$, chunk size $k=4$). $\uparrow$ indicates higher is better, while $\downarrow$ indicates lower is better.

SCR provides the best trade-off between quality and speed

EXPERIMENTS

Analysis 2: Controllability (α -Tuning)

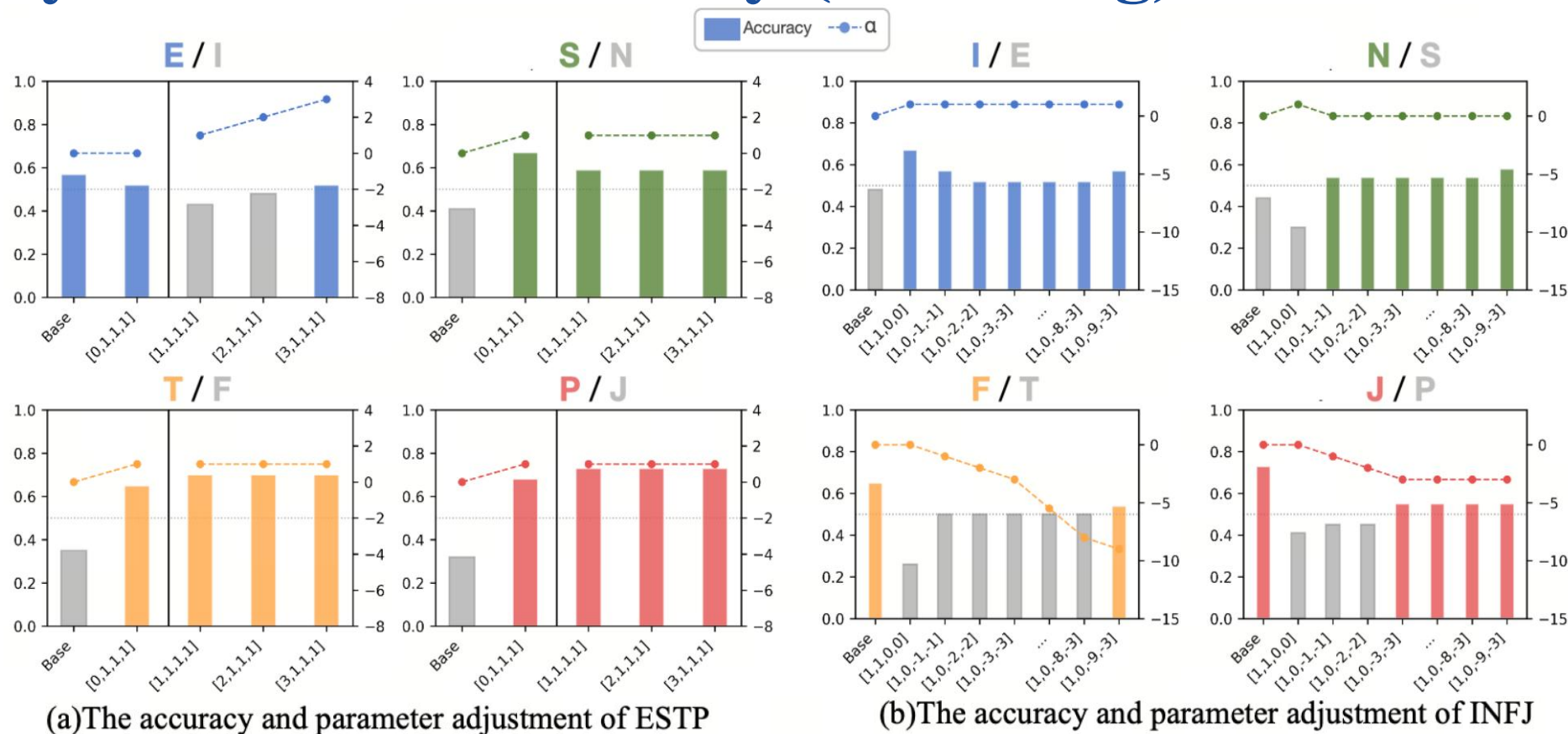


Figure: Iterative tuning process for the α . Bars indicate prediction accuracy (left axis) for each MBTI dimension; dashed lines track α values (right axis) at each optimization step.



CIKM SEOUL 2025
November 10 - 14



Thank you for your attention!

Paper



WeChat



Email: chenrongxin24s@ict.ac.cn