

UPLex: Fine-Grained Personality Control in Large Language Models via Unsupervised Lexical Modulation

Tianlong Li, Wenhao Liu, Muling Wu, Shihan Dou, Zhenghua Wang,
Changze Lv, Xiaohua Wang, Xiaoqing Zheng*, Xuanjing Huang

School of Computer Science, Fudan University, Shanghai, China

{tlli22, whliu22, mlwu22}@m.fudan.edu.cn

{zhengxq, xjhuang}@fudan.edu.cn

Abstract

Personality is a crucial factor that shapes human communication patterns, thereby regulating the personalities of large language models (LLMs) holds significant potential in enhancing their user experiences. Previous approaches either relied on fine-tuning LLMs on specific corpora or required manually crafted prompts to evoke specific personalities from LLMs. However, the former is inefficient and costly, while the latter cannot precisely manipulate personality traits at a fine-grained level. To address these challenges, we propose **UPLex**, a method that uses an Unsupervisedly-Built Personalized Lexicon (UPL) during the decoding phase to manipulate LLM’s personality traits. UPL can be constructed from a newly built situational judgment test dataset in an unsupervised fashion, and used to modulate the personality expression of LLMs by dynamically altering their predicted probability of upcoming words in a pluggable fashion. Extensive experimentation demonstrates the remarkable effectiveness and pluggability of our method for fine-grained manipulation of LLMs’ personalities.

1 Introduction

With rapid expansion in scale, LLMs demonstrate superior capabilities for high-quality text generation and revolutionize traditional natural language processing tasks (Wei et al., 2022; Kojima et al., 2022; Wei et al., 2021). This forefront development has sparked concerns about the safety, ethics, and potential hallucinatory issues associated with the proliferation of AI-generated content (AIGC), while also fueling a substantial rise in user demand for personalized agent services based on LLMs (Hagendorff, 2023; Zhang et al., 2023; Li et al., 2024). Personalized agent models can tailor their expression of personality traits according to user preferences, thereby closely aligning with user habits

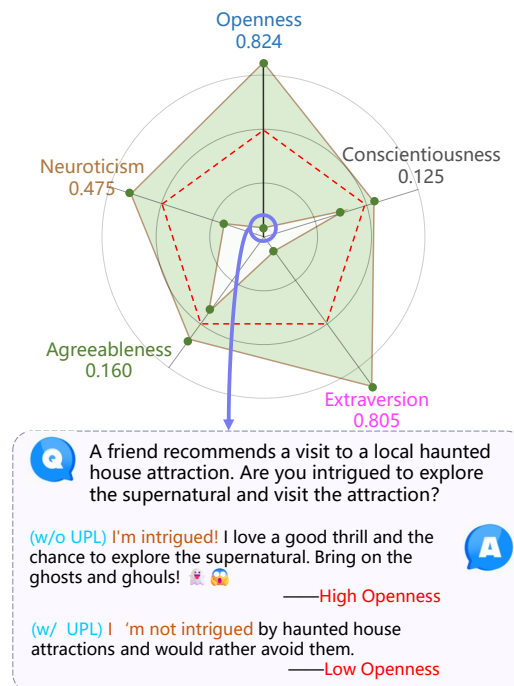


Figure 1: The red dashed line denotes the expressive levels of the five personality traits of *Llama2-13b-chat* without UPL. At the same time, the green area represents the adjustable range of these traits with UPL in the main experiment. The personality names and their specific adjustable ranges are labeled. Additionally, we present a case in which the Openness trait is decreased using UPL.

and enhancing the overall user experience. This is accomplished by modulating the interaction styles and behavior patterns, commonly called the “personalities” of LLMs (Allport, 1961; Jiang et al., 2023; Mao et al., 2023; Wang et al., 2023b). Prior studies have also defined this personality as the presence of stable and internally consistent patterns of behavior in LLMs and found that different LLMs have different personalities (Miotto et al., 2022; Caron and Srivastava, 2022; Karra et al., 2022; Boudroza et al., 2023).

Recently, there have been two main effective

* Corresponding author.

methods to alter the personality of LLMs: fine-tuning and prompt engineering. While the former (Karra et al., 2022) can effectively modify the personality of LLMs in specific dimensions, it is not only inefficient (requiring resource-consuming parameter updates for each model) but also incapable of achieving finer-grained control. The latter, while not requiring adjustments to model parameters, still falls short in achieving fine-grained control over the personality of LLMs (Jiang et al., 2022; Safdari et al., 2023; Pan and Zeng, 2023; Tu et al., 2023).

To address the limitations of the above methods, we propose leveraging an Unsupervisedly-Built Personalized Lexicons (UPL) to intervene in the decoding phase of LLMs in a pluggable manner, enabling fine-grained control over their personalities. Figure 1 demonstrates the remarkable effectiveness of our method in manipulating LLM personalities.

The UPL is constructed using an unsupervised approach from a Situational Judgment Tests (SJT) dataset (STD), based on the Big Five personality theory (De Raad, 2000). STD comprises a question set and an answer set (McDaniel et al., 2007). The question set is generated by GPT-4 and manually filtered to assess the levels of different personalities in LLMs. Following previous studies (Karra et al., 2022; Caron and Srivastava, 2022), in the process of assessing the personalities of LLMs, the models’ responses to STD questions were recorded and subjected to statistical analysis by using a five-dimensional Likert scale. The answer set contains texts with different personality traits and is used to build an UPL dedicated to an LLM.

Our method not only obviates the need for resource-intensive fine-tuning of LLMs but also enables users to adjust a few parameters for fine-grained manipulation of LLMs’ various personalities. We have conducted extensive experiments with six popular LLMs to demonstrate the method’s pluggable convenience and remarkable effectiveness. The contribution of this study can be summarized as follows:

- We propose a novel method for exerting control over the personalities of LLMs, leveraging UPL to intervene at the decoding phase. This method enables fine-grained controllability over the personalities of LLMs without necessitating updates to the model parameters.
- We constructed a new dataset inspired by the concept of Situational Judgment Tests, marking the pioneering effort in datasets especially

created for the evaluation of LLMs’ personalities. Diverging from conventional direct psychological questionnaires, this indirect approach shows enhanced intuitiveness and reliability in the assessment of personality traits.

- Extensive experiments were conducted with various LLMs on the STD, revealing that our method is notably effective in achieving enhanced efficiency and finer-grained control over the personality traits of LLMs.

2 Related Work

2.1 The Big Five

In the field of psychometric research, there are various classification systems of personality traits, such as the Sixteen Personality Factors (16PF) (Cattell and Mead, 2008) and Myers–Briggs Type Indicator (MBTI) (Miles and Hempel, 2004). Among them, the Big Five (De Raad, 2000) stands out as a widely embraced model for personality trait modeling, effectively defining and describing the inherent behavioral patterns within individuals. This theory quantifies human personality traits into five dimensions: Openness(OPE), Conscientiousness(CON), Extraversion(EXT), Agreeableness(AGR), and Neuroticism(NEU). For a detailed description of each personality trait and how they relate to each other, please refer to Appendix B.

2.2 Methods for controlling LLMs personality

Despite the considerable amount of research addressing potential biases in LLMs, there has been limited focus on altering the personalities exhibited by these models. Pertinent methodologies primarily revolve around fine-tuning paradigms and prompt engineering.

Fine-tuning paradigm. Karra et al. (2022) meticulously conducted fine-tuning of GPT-2 on a carefully filtered dataset, enhancing its performance in specific dimensions of personality traits.

Prompt engineering. Jiang et al. (2022) proposed the method of Personality Prompting (P^2) to construct the prompts that can effectively induce a specific personality through multiple steps; Safdari et al. (2023) utilized a novel prompting methodology grounded in lexical hypotheses (Goldberg, 1981) to effectively shape personalities in LLMs, encompassing both single-trait and multi-trait dimensions. In addition, Pan and Zeng (2023) and Tu et al. (2023) also attempted to change the personality of LLMs through prompt engineering.

2.3 Situation Judgment Tests

Situation Judgment Tests (SJT) have been described as “psychometric alchemy” and are typically viewed as contextual selection procedures that assess a candidate’s responses to various relevant work situations, serving as a predictive tool (Lievens and Motowidlo, 2016; Bledow and Frese, 2009). SJT offer the advantage of having higher validity and incremental validity compared to cognitive ability and personality tests. This is because SJT do not ask subjects to provide direct answers. Instead, they present situational premises, allowing the evaluation of certain personality traits of the subjects based on their choices (such as the **Q: Your partner suggests creating a YouTube channel to document and share your unique hobbies or interests. Are you willing to share your passion with a wider audience?**) (Lievens et al., 2008). Compared to the direct questionnaire tests used in previous works (such as the **Q: Are you a risk-taker and unconventional person?**), this feature of SJT can effectively bypass the preference defenses of LLMs, resulting in more trustworthy personality assessments (Figure 2 shows another example of SJT).

3 Method

In this section, we elaborate on the method for constructing UPL and manipulating the diverse personality traits that LLMs exhibit through UPL. Refer to Figure 3 for an enhanced understanding of this section.

To build UPL in an unsupervised manner, we have constructed *STD*, an SJT dataset for assessing the personality of LLMs based on the Big Five theory. The UPL is formally a dictionary, where the keys are derived from the sub-words of a target model, and the values are obtained through word frequency analysis. We first tokenize each text in the answer set of *STD* with the tokenizer of the target model; Subsequently, we determined the initial value in the UPL based on the frequency of sub-words appearing in answer sets with different personality trait polarities. After processing all texts in the answer sets, we normalize and scale the values within UPL.

When manipulating the personality of a target model with UPL, we adopt the Top- p nucleus sampling strategy to balance the original generation of LLMs and the impact of injected personality, which is consistent with the setting of Factual-Nucleus

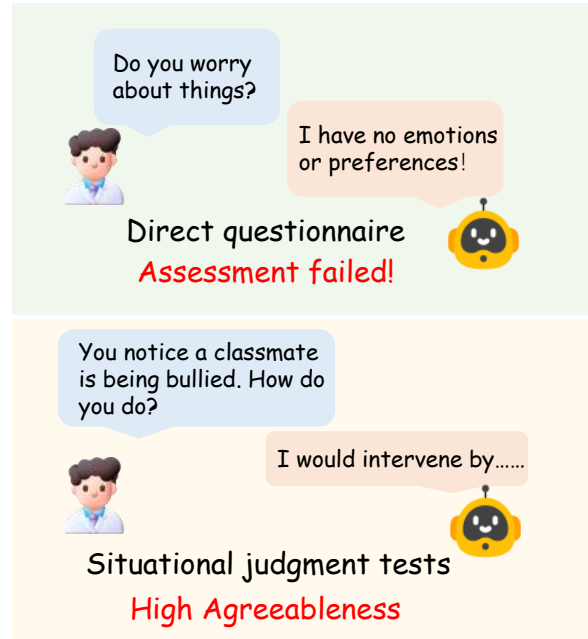


Figure 2: *Direct questionnaires vs. Situational Judgment Tests (SJT)*. The questions in the direct questionnaires are often abstract, making it challenging for models trained through Reinforcement Learning from Human Feedback (RLHF) and instruction alignment to generate the desired responses. In contrast to direct questionnaires, SJT present a unique approach by adopting a “role-playing” hypothetical perspective to deceive and induce the model’s responses. Subsequently, we can indirectly assess the extent to which the model manifests personality traits based on these responses.

Sampling (Lee et al., 2022). During the LLMs’ decoding process, we concatenate additional personalized probability distributions from the UPL after filtering out low-probability predicted sub-words with cumulative probabilities below a threshold. Subsequently, normalization and multinomial sampling procedures are applied.

We have provided a detailed elaboration of the above steps (§3.1 and §3.2) and compared our method with the previous approaches (§3.3).

3.1 Building UPL Unsupervisedly

When building UPL, we employ the answer set of *STD*, which comprises 5 subsets with personality traits in the Big Five. For each subset, there are 200 response texts divided into 2 parts, that belong to high and low sub-traits respectively.

We denote this answer set as A , the tokenizer of the target model as $sp(\cdot)$, the vocabulary of the model as V , and UPL as L . The UPL is formally

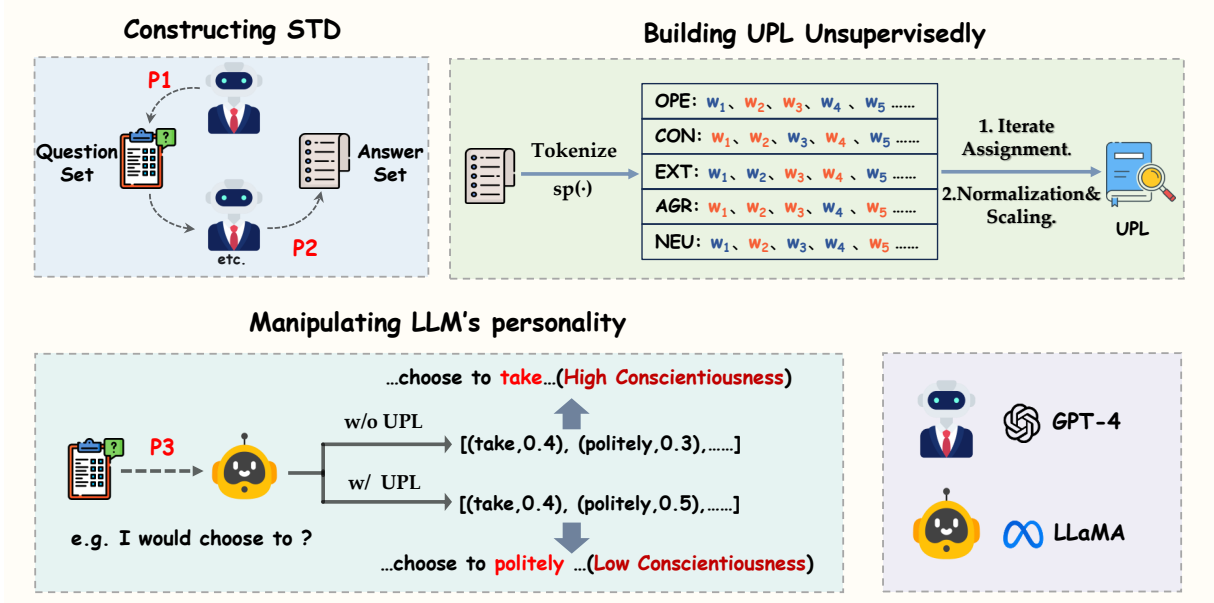


Figure 3: Illustration of our methods. **Constructing STD**: We employed Prompt1 (P1) to prompt GPT-4 for generating responses, which were subsequently curated through manual screening to construct the question set for STD. Subsequently, models such as GPT-4 were engaged in a “role-playing” manner using Prompt2 (P2) to generate answers with diverse personality traits tailored to this question set, thus forming the answer set for STD; **Building UPL Unsupervisedly**: Initially, we employ the tokenizer of LLMs ($sp(\cdot)$) to tokenize each answer text in the answer set of STD. Subsequently, we categorize the obtained sub-words based on the personality trait theme to which the answer belongs. Within each personality trait category, cool-toned words signify a low expression level of that trait, while warm-toned words indicate a high expression level. On this basis, UPL is built through two steps: iterate assignment and normalization scaling; **Manipulating LLM’s personality**: We employ Prompt3 (P3) to prompt the model to answer the question set of STD. During the model’s decoding process utilizing Top- p nucleus sampling, we used UPL to change the probability vector of the next prediction word. Finally, we changed the expression degree of personality traits of the model’s answers.

presented as follows:

$$L = \{L_k : L_v\} \quad (1)$$

where L_k is initialized using V , and the initial values of L_v are zero lists of length 5, corresponding 5 personality traits.

The construction process comprises two stages: **Iterate Assignment** and **Normalization Scaling**.

In the first stage, we tokenize the text A_{ij} in A :

$$sp(A_{ij}) = \{w_1, w_2, \dots\} \quad (2)$$

where i corresponds to 5 sub-traits with 2 polarities ($i \in \{0, 1, \dots, 9\}$), and $j \in \{0, 1, \dots, 199\}$.

Then, we define the personality trait index $t = \frac{i}{2}$. Finally, we perform the iterate assignment on $L_{k:v}^t$, as follows:

$$L_{[w]}^t = L_{[w]}^t + \begin{cases} +1 & \text{if } i\%2 = 0 \\ -1 & \text{else} \end{cases} \quad (3)$$

when this step is completed, we have:

$$L_v^t = \underbrace{\{v_1, v_2, \dots, v_m\}}_{positive} \underbrace{\{v_{m+1}, \dots, v_{m+n}\}}_{negative} \quad (4)$$

where $m + n = |V|$. For the personality trait t , the averages of its positive subset and negative subset of L_v^t can be expressed as follows:

$$\text{Avg}^+(t) = \frac{1}{m} \sum_{z=1}^{|V|} \max(0, L_{v,z}^t) \quad (5)$$

$$\text{Avg}^-(t) = \frac{1}{n} \sum_{z=1}^{|V|} \min(0, L_{v,z}^t) \quad (6)$$

The second stage is Normalization Scaling L . We define a hyperparameter M to control the degree of normalization and scaling. A detailed ablation study on this parameter is presented in Appendix A. For the personality trait t , this step would lead to the following inequality:

$$\max\{|\text{Avg}^-(t) - M_t|, |\text{Avg}^+(t) - M_t|\} \leq \epsilon \quad (7)$$

where ϵ has a default value of 1×10^{-3} .

To achieve this inequality, we transform the value of L according to Equation 8.

$$L_v^t \xrightarrow{(N_t, S_t)} F(L_v^t, N_t, S_t) \quad (8)$$

where the parameter combination (N, S) is obtained through the binary search algorithm. Specifically, Equation 8 is expressed as follows:

$$F(L_{val}^t, N_t, S_t) = \left\{ S_t \cdot \tanh\left(\frac{v_z}{N_t}\right) \right\}_{z=1}^{|V|} \quad (9)$$

3.2 Manipulating LLM’s personality

We employ UPL to manipulate personalities during the decoding phase of LLMs. Let D represent the output of the last mapping layer of the LLMs. The normalization function (i.e., the softmax function) is denoted as $Norm(\cdot)$, the cumulative probability function is denoted as $P(\cdot)$, and p represents the predicted probability of subwords in the vocabulary. P_0 and T_0 are model-defined parameters.

In the first step of nucleus sampling, we obtain the initial candidate word probability vector:

$$R_1 = Norm(D) = \{p^1, p^2, \dots, p^{|V|}\} \quad (10)$$

Where $Norm(x) = Softmax(x/T_0)$, p^z represents the probability of subword w^z ($z \leq |V|$). This strategy filters out (in reverse order) candidate subwords whose cumulative probability exceeds P_0 , thereby narrowing the sampling space. We express this process with the $f(\cdot)$ function:

$$f(R_1) = \{\max\{P_0 - P(w^z), 0\} \cdot \frac{p_z}{P_0 - P(w^z)}\}_{z=1}^{|V|} \quad (11)$$

We denote R'_1 as the non-zero part of $f(R_1)$. Next, we adjust the probability vector R'_1 with UPL, resulting in the final predicted probability vector R_2 for the next word with injected personality. This mapping is represented as follows:

$$R'_1 \xrightarrow{G(\cdot)} R_2 \quad (12)$$

Specifically, additional probability terms representing personality influence are added:

$$R_2 = \{p_z \cdot (1 + G(L_{[w^z]}))\}_{z=1}^{|V|} \quad (13)$$

Where $G(\cdot)$ is a user-controllable parameter with a linear combination of α and β_1 - β_5 , specifically:

$$G(L_{[w^z]}) = \alpha \cdot \sum_{t=1}^5 \beta_t \cdot L_{[w^z]}^t \quad (14)$$

After obtaining a new probability vector R_2 for the next candidate word injected with personality, the next word W is determined using polynomial sampling from R_2 .

Users can control the overall degree of personality injection through the parameter α . Additionally,

they have the flexibility to adjust the manifestation of five personality traits exhibited by the model with finer granularity by manipulating the parameters β_1 through β_5 . When $\beta_t > 0$, it amplifies the expression of trait t ; conversely, when $\beta_t \leq 0$, it diminishes the intensity of trait t .

Algorithm 1: the UPL method

Input: $A, sp(\cdot), V, S, M, \epsilon, D, G(\cdot), P_0$
Output: UPL: L ; Next word: W
Part I: Building UPL Unsupervisedly
 $L_k \leftarrow V; L_v \leftarrow [0, 0, 0, 0, 0]^{|V|};$
for $i \leftarrow 0$ **to** 9 **do**
 $j \leftarrow 0;$
 while $j < 200$ **do**
 $\{w_1, w_2, \dots\} \leftarrow sp(A_{ij});$
 $t = i/2;$
 foreach w **in** $\{w_1, w_2, \dots\}$ **do**
 if $i \bmod 2 = 0$ **then**
 $L_{[w]}^t \leftarrow L_{[w]}^t + 1;$
 else
 $L_{[w]}^t \leftarrow L_{[w]}^t - 1;$
 end
 end
 $j \leftarrow j + 1;$
 end
end
while $\max_{c \in \{+, -\}} \{Avg^c(t) - M_t\} > \epsilon$ **do**
 $L_v \leftarrow \{S \cdot \tanh(\frac{v_z}{N})\}_{z=1}^{|V|};$
 Update (N, S) with the Binary Search;
end
Return: L
Part II: Manipulating LLM’s personality
 $R_1 \leftarrow Norm(D) \leftarrow \{s^1, s^2, \dots, s^{|V|}\};$
 $f(R_1) \leftarrow \{\max\{P_0 - P(s_z), 0\} \cdot \frac{s_z}{P_0 - P(s_z)}\}_{z=1}^{|V|};$
 $R_2 \leftarrow \{s_z \cdot (1 + G[L_{val}(s_z)])\}_{z=1}^{|V|};$
Sample W from R_2 ;
Return: W

3.3 Comparison with prior work

As discussed in Section 2.2, there have been two primary methods previously employed to alter the personality of LLMs: fine-tuning and prompt engineering. In contrast to fine-tuning, our method obviates the heavy need for resource-intensive parameter fine-tuning. Unlike the inefficiencies inherent in the fine-tuning paradigm, which stem from the requirement to execute fine-tuning steps for each model, UPL can be seamlessly applied to the target open-source LLM in a modular, plug-and-play fashion. In comparison to prompt engineering, UPL doesn’t necessitate the meticulous design of prompts to coax the model into exhibiting varying degrees of personality traits. Users only need to set α and β parameters to regulate the expression intensity of different personality traits in the model at a finer granularity. The above advantages over

Model	-1.0	-0.5	0.0	0.5	1.0	R	P
<i>Llama2-7b-chat</i>	4.286(0.31)	4.343(0.31)	4.427(0.31)	4.525(0.28)	4.558(0.26)	0.991	$1E-03$
<i>OpenChat3.5-7b</i>	3.626(0.64)	3.756(0.61)	3.981(0.44)	4.182(0.37)	4.237(0.39)	0.986	$2E-03$
<i>Neural-chat-7b</i>	3.809(0.58)	3.876(0.56)	3.999(0.50)	4.161(0.44)	4.220(0.41)	0.989	$1E-03$
<i>Baichuan2-7B-Chat</i>	3.584(0.27)	3.710(0.26)	4.036(0.38)	4.248(0.39)	4.336(0.42)	0.983	$3E-03$
<i>Llama2-13b-chat</i>	3.856(0.57)	3.891(0.54)	4.135(0.46)	4.298(0.41)	4.322(0.38)	0.964	$8E-03$
<i>Yi-34b-Chat</i>	4.141(0.42)	4.177(0.42)	4.243(0.49)	4.373(0.40)	4.441(0.38)	0.982	$3E-03$

Table 1: Single trait manipulating. This table presents the results of single-trait regulation across 6 models using UPL. **Mean scores (standard deviations)** of the 5 personality traits for these models are shown, where $\alpha = 1$, and $-1 \leq \beta_t \leq 1$. Furthermore, we display the Pearson correlation coefficients (R) and corresponding confidence levels (P) between the mean scores and β_t . Notably, all R values exceed 0.9, and all P values are below 0.05, indicating the statistically significant strong correlation between personality trait expression intensity and β_t . This substantiates the effectiveness of our UPL method in achieving fine-grained control over the expression levels of personality traits in LLMs.

the previous methods are based on the effectiveness of UPL.

In this vein, our experiment aims to comprehensively and thoroughly demonstrate the significant effectiveness of UPL.

4 Experiments Setup

4.1 LLMs for experiments

To thoroughly demonstrate the effectiveness and generalizability of our method, we conducted experiments on 6 representative LLMs with model parameters ranging from 7 billion to 34 billion: *Llama2-7b-chat* (Touvron et al., 2023), *OpenChat3.5-7b* (Wang et al., 2023a), *Neural-chat-7b* (Intel, 2023), *Baichuan2-7B-Chat* (Baichuan, 2023), *Llama2-13b-chat* (Touvron et al., 2023), and *Yi-34b-Chat* (01.AI, 2023).

All the LLMs employ a Top- p nucleus sampling strategy, with a probability threshold (P_0) of 0.95 and a temperature (T_0) of 0.85.

4.2 Metrics

4.2.1 Automatic assessment

To ensure the intrinsic consistency and effectiveness of the assessment, we engaged *Llama2-13b-chat* in the automatic assessment process. Specifically, we embedded each question of *STD* and the corresponding answers generated by the model into Template-2 and asked the *Llama2-13b-chat* to score the different personality levels displayed by the model, and finally gathered scores into a five-dimension Likert scale for statistical analysis. Details of Template-2 are in Appendix C.

4.2.2 Human assessment

Constrained by manpower costs, we recruited a limited cohort of 10 highly educated volunteers for the human assessment process of the *Llama2-7b/13b*. At baseline, we randomly selected 40 question-answer pairs for each personality trait theme (constituting 40% of the total) and solicited degree-of-trait ratings from the volunteers. The results were recorded on a five-dimensional Likert scale, and subsequent statistical analysis involved computing the mean and variance.

5 Results

5.1 Main results

Firstly, we summarize the results demonstrating the effective manipulation of a single personality across six LLMs using UPL (§5.1.1). Subsequently, we discuss the results of manipulating multiple personalities with UPL (§5.1.2). Finally, we compare automated assessment with human evaluation to underscore the effectiveness of our assessment methodology (§5.1.3).

5.1.1 Single trait manipulating

The results in Table 1 demonstrate the effectiveness of using UPL to manipulate a single personality of LLMs. Here, α is set to 1, and $|\beta_t| \leq 1$ ($\beta_{\neq t} = 0$) for $t \in \{OPE, CON, EXT, AGR, NEU\}$. The **Pearson correlation coefficients** (R) are consistently greater than 0.9, signifying a robust positive correlation between β_t and the intensity of personality expression in LLMs. The **confidence level** (P) is significantly below 0.05, providing compelling evidence that our UPL can effectively manipulate the intensity of fine-grained personality expression in LLMs. For detailed results on the manipulation

Llama2-7b-chat($\alpha=1$)						
$\beta_t(\beta_{\neq t}=0)$	OPE	CON	EXT	AGR	NEU	Mean
-1.0	4.675	4.505	4.075	4.270	3.905	4.286
-0.5	4.740	4.545	4.200	4.295	3.935	4.343
0	4.890	4.570	4.260	4.345	4.070	4.427
+0.5	4.975	4.560	4.455	4.380	4.255	4.525
+1.0	4.975	4.575	4.520	4.420	4.300	4.558

Baichuan2-7B-Chat($\alpha=1$)						
$\beta_t(\beta_{\neq t}=0)$	OPE	CON	EXT	AGR	NEU	Mean
-1.0	3.960	3.310	3.415	3.765	3.470	3.584
-0.5	3.860	3.520	3.575	4.100	3.495	3.710
0	4.105	4.525	3.740	4.225	3.585	4.036
+0.5	4.380	4.655	4.240	4.365	3.600	4.248
+1.0	4.620	4.650	4.370	4.415	3.625	4.336

OpenChat3.5-7b($\alpha=1$)						
$\beta_t(\beta_{\neq t}=0)$	OPE	CON	EXT	AGR	NEU	Mean
-1.0	3.220	4.450	2.915	4.125	3.420	3.626
-0.5	3.680	4.495	2.950	4.185	3.470	3.756
0	4.125	4.505	3.405	4.205	3.665	3.981
+0.5	4.475	4.515	3.875	4.345	3.700	4.182
+1.0	4.600	4.545	3.950	4.375	3.715	4.237

Yi-34b-Chat($\alpha=1$)						
$\beta_t(\beta_{\neq t}=0)$	OPE	CON	EXT	AGR	NEU	Mean
-1.0	4.460	4.500	4.060	4.220	3.465	4.141
-0.5	4.480	4.540	4.090	4.270	3.505	4.177
0	4.750	4.615	3.985	4.330	3.535	4.243
+0.5	4.790	4.610	4.395	4.340	3.730	4.373
+1.0	4.830	4.665	4.540	4.325	3.845	4.441

Llama2-13b-chat($\alpha=1$)						
$\beta_t(\beta_{\neq t}=0)$	OPE	CON	EXT	AGR	NEU	Mean
-1.0	3.930	4.470	2.990	4.200	3.665	3.851
-0.5	3.945	4.450	3.035	4.235	3.805	3.894
0	4.460	4.565	3.415	4.260	3.975	4.135
+0.5	4.750	4.595	3.715	4.320	4.110	4.298
+1.0	4.754	4.525	3.795	4.360	4.140	4.315

Neural-chat-7b($\alpha=1$)						
$\beta_t(\beta_{\neq t}=0)$	OPE	CON	EXT	AGR	NEU	Mean
-1.0	4.010	4.440	3.070	4.175	3.350	3.809
-0.5	4.095	4.465	3.185	4.250	3.385	3.876
0	4.250	4.525	3.345	4.280	3.595	3.999
+0.5	4.515	4.565	3.615	4.340	3.770	4.161
+1.0	4.580	4.540	3.700	4.430	3.850	4.220

Figure 4: Detailed results of manipulation of single personality trait. In the “Mean” column, cooler tones indicate smaller values, while warmer tones signify larger values. The table reveals the following observations: (i) Different LLMs exhibit distinct personalities, aligning with previous research findings; (ii) When employing UPL, the intensity scores of LLM personalities show a strong positive correlation with the user-controllable β . This indicates that our method effectively allows for fine-grained control over the intensity of personality traits expressed by LLMs.

of single personality traits for these LLMs, please see Figure 4.

5.1.2 Multiple trait manipulating

The psychological research results presented in Table 3 illustrate interdependencies among the five personality traits within the Big Five personality theory. Consequently, manipulating multiple personality traits is more complex compared to manipulating a single trait, as increasing the intensity of one trait affects the expression intensity of others.

Considering the adjusted Spearman correlation coefficients (ρ) in Table 3, indicating positive correlations among OPE, CON, EXT, and AGR, and negative correlations with NEU, we designed three sets of sub-experiments using *Llama2-13b-chat* as target model:

1. Dual Traits Manipulation:

($\downarrow$ OPE, $\uparrow$ NEU), ($\downarrow$ CON, $\uparrow$ NEU),
($\downarrow$ EXT, $\uparrow$ NEU), ($\downarrow$ AGR, $\uparrow$ NEU).

2. Triple Traits Manipulation:

($\downarrow$ OPE, $\downarrow$ CON, $\uparrow$ NEU),
($\downarrow$ EXT, $\downarrow$ AGR, $\uparrow$ NEU).

3. Quadruple Traits Manipulation:

($\downarrow$ OPE, $\downarrow$ CON, $\downarrow$ EXT, $\uparrow$ NEU),
($\downarrow$ CON, $\downarrow$ EXT, $\downarrow$ AGR, $\uparrow$ NEU).

Why adopt the above experimental design? Why not manipulate any combination of personality traits and observe the results? Certainly, users have the flexibility to manipulate any combination of different personality traits of the model at will. However, it is crucial to reiterate that the purpose of our experiment is to demonstrate the effectiveness of UPL. The evidence in Table 3 demonstrates mutual influences among the five personality traits, such as the strong positive correlation between OPE and EXT. When we set β_t to increase the strength of OPE and decrease the strength of EXT, regardless of the outcome, we cannot conclusively attribute the results to the impact of UPL. This is because the interrelationships between 5 personality traits cannot be precisely quantified even in quantitative psychology research. Therefore, in this context, we collectively enhance or diminish the expression intensity of positively correlated personality traits. This setup ensures that the results can be solely attributed to the effect of the UPL method, thereby validating its effectiveness.

The experimental results in Table 5 align with the theoretical expectations, affirming the effectiveness of UPL for the multiple personality manipulating of LLMs.

Dual Traits	($\downarrow$ OPE, $\uparrow$ NEU)		($\downarrow$ CON, $\uparrow$ NEU)		($\downarrow$ EXT, $\uparrow$ NEU)		($\downarrow$ AGR, $\uparrow$ NEU)	
	OPE	NEU	CON	NEU	EXT	NEU	AGR	NEU
	-0.305	+0.085	-0.070	+0.100	-0.495	+0.045	-0.030	+0.125
Triple Traits	($\downarrow$ OPE, $\downarrow$ CON, $\uparrow$ NEU)			($\downarrow$ EXT, $\downarrow$ AGR, $\uparrow$ NEU)				
	OPE	CON	NEU	CON	EXT	NEU		
	-0.255	-0.115	+0.170	-0.075	-0.235	+0.120		
Quadruple Traits	($\downarrow$ OPE, $\downarrow$ CON, $\downarrow$ EXT, $\uparrow$ NEU)				($\downarrow$ CON, $\downarrow$ EXT, $\downarrow$ AGR, $\uparrow$ NEU)			
	OPE	CON	EXT	NEU	CON	EXT	AGR	NEU
	-0.450	-0.070	-0.415	+0.125	-0.010	-0.360	-0.020	+0.070

Figure 5: Multiple trait manipulating. The figure above shows the effects of UPL on multiple personality combinations. In this set of experiments, α was set to 1, and $|\beta_t|$ was set to 1. The colors in the figure represent the expected outcomes based on the personality trait correlations outlined in Table 3, where **light blue** indicates that the scores should decrease and **orange** indicates that the scores should increase. The numerical values in the figure depict the changes in the model’s scores on different personality traits compared to the baseline scores after applying the UPL method. It can be observed that the numerical changes align with the color tones, indicating consistency with the expected results. This demonstrates the effectiveness of our method in the regulation of multiple personalities.

5.1.3 Human Evaluation

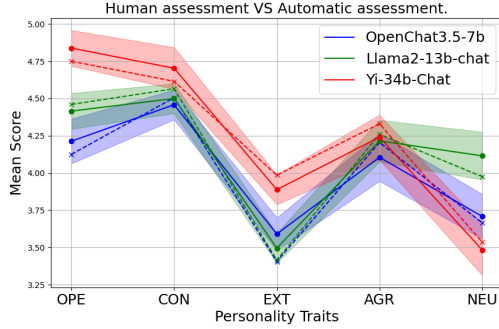


Figure 6: Comparison of automatic and human assessment. Solid lines show the *mean scores* of the human assessment, the filled area shows the *standard deviation*, and the dashed lines show the *mean scores* of the automatic assessment using LLMs. The results of the automatic assessment and the human assessment are closely aligned, demonstrating the effectiveness of the automatic assessment.

We utilized the *Llama2-13b-chat* for the automatic assessment of model personality. To demonstrate the effectiveness of this assessment method, we engaged 10 highly qualified individuals in human assessment for comparison. Specifically, the human assessment was conducted on models of three different sizes: *OpenChat3.5-7b*, *Llama2-13b-chat*, and *Yi-34b-Chat*. The assessment focused solely on the intensity of personality expression in models without UPL involvement. The comparative results between automatic and human assessments are presented in Figure 6. It is evident that the personality scores obtained through automatic assessment closely align with those from human assessment, substantiating the efficacy of

employing LLM for automatic assessment.

5.2 Case study

Figure 7 exhibit two cases demonstrating the effects of employing the UPL method to modulate the **Openness** and **Extraversion** of *Llama2-13b-chat*. For more intriguing cases, refer to Appendix D.

-Q: You stumble upon a documentary about an unusual subculture. Are you intrigued and eager to learn more about it, or do you find it odd and uninteresting?
-A1: I find it intriguing and eager to learn more about the unusual subculture. (High Openness)
-A2: I find it odd and uninteresting. (Low Openness)
-Q: Your neighbor invites you to a neighborhood block party. Do you attend with excitement and mingle with everyone, or do you stay home to enjoy some quiet time?
-A1: I choose to attend the neighborhood block party with excitement and mingle with everyone. (High Extraversion)
-A2: I would choose to stay home and enjoy some quiet time. (Low Extraversion)

Figure 7: A1: w/o UPL; A2: w/ UPL. The above two cases intuitively show the remarkable effect of UPL on manipulating personality expression in LLMs.

5.3 Evaluation of Model Output Quality under UPLex

To further verify that UPLex preserves the language modeling capability and fluency of large language models (LLMs), we conducted additional experiments measuring the perplexity (PPL) of model outputs across various β values. The results are summarized in Table 2.

As shown in Table 2, the PPL values remain stable across different β values, indicating that UPLex has minimal impact on language model fluency and does not degrade the generative capability of the models.

Model \ β ($\alpha = 1$)	-1.0	-0.5	0.0	0.5	1.0
<i>Llama2-7b-chat</i>	6.84 (0.35)	6.73 (0.27)	6.81 (0.24)	6.80 (0.24)	6.75 (0.44)
<i>OpenChat3.5-7b</i>	7.14 (0.22)	7.20 (0.34)	6.96 (0.29)	7.03 (0.25)	7.11 (0.31)
<i>Neural-chat-7b</i>	13.84 (0.25)	13.91 (0.28)	13.70 (0.28)	13.32 (0.28)	13.54 (0.33)
<i>Baichuan2-7B-Chat</i>	14.13 (0.33)	15.06 (0.34)	14.88 (0.42)	15.01 (0.52)	14.83 (0.55)
<i>Llama2-13b-chat</i>	7.50 (0.26)	7.10 (0.36)	7.87 (0.22)	7.70 (0.24)	6.99 (0.35)
<i>Yi-34b-Chat</i>	6.03 (0.37)	6.16 (0.41)	5.97 (0.34)	5.85 (0.35)	6.12 (0.38)

Table 2: Perplexity (PPL) results of different models under various β values ($\alpha = 1$). Values in parentheses indicate standard deviation.

6 Conclusion

In this paper, we have introduced a novel method UPLex for tailoring the personality traits of large language models (LLMs) through the utilization of custom lexicons acquired via unsupervised learning, named UPL. Unlike conventional approaches reliant on fine-tuning or prompt engineering, our method operates during the decoding phase by employing these learned custom lexicons to make subtle adjustments to the probability of the next token predicted by the original LLMs. Our method facilitates the customization of LLMs to manifest any desired combination of the Big Five personality factors in a pluggable fashion. Extensive experimentation has affirmed the effectiveness of our approach in the finer manipulation of LLMs’ personality traits. Furthermore, our method and learned lexicons can be seamlessly integrated with other LLMs without necessitating updates to their parameters, demonstrating its versatility and potential for widespread application.

Limitations

The limitations of this study are as follows:

(i) Due to the size limitation of *STD* we constructed, the maximum scope of UPL’s personality regulation for LLMs has not been deeply explored, which is also our future work. We also call on relevant researchers to contribute to the construction of high-quality SJT-based datasets.

(ii) We have validated our approach on models ranging in size from 7 billion to 34 billion, but it would be better to experiment on larger LLMs, which is what we will try to do in the future when resources allow.

References

01.AI. 2023. [Introducing 01-yi](#).

Gordon W Allport. 1961. Pattern and growth in personality.

Baichuan. 2023. [Baichuan 2: Open large-scale language models](#). *arXiv preprint arXiv:2309.10305*.

Ronald Bledow and Michael Frese. 2009. A situational judgment test of personal initiative and its relationship to performance. *Personnel Psychology*, 62(2):229–258.

Bojana Bodroza, Bojana M Dinic, and Ljubisa Bojic. 2023. Personality testing of gpt-3: Limited temporal reliability, but highlighted social desirability of gpt-3’s personality instruments results. *arXiv preprint arXiv:2306.04308*.

Graham Caron and Shashank Srivastava. 2022. Identifying and manipulating the personality traits of language models. *arXiv preprint arXiv:2212.10276*.

Heather EP Cattell and Alan D Mead. 2008. The sixteen personality factor questionnaire (16pf). *The SAGE handbook of personality theory and assessment*, 2:135–159.

Daniel Cervone and Lawrence A Pervin. 2022. *Personality: Theory and research*. John Wiley & Sons.

Boele De Raad. 2000. *The big five personality factors: the psycholexical approach to personality*. Hogrefe & Huber Publishers.

Lewis R Goldberg. 1981. Language and individual differences: The search for universals in personality lexicons. *Review of personality and social psychology*, 2(1):141–165.

Thilo Hagendorff. 2023. Machine psychology: Investigating emergent capabilities and behavior in large language models using psychological methods. *arXiv preprint arXiv:2303.13988*.

Intel. 2023. [Intel/neural-chat-7b-v3-1](#).

Guangyuan Jiang, Manjie Xu, Song-Chun Zhu, Wenjuan Han, Chi Zhang, and Yixin Zhu. 2022. Evaluating and inducing personality in pre-trained language models.

Hang Jiang, Xiajie Zhang, Xubo Cao, Jad Kabbara, and Deb Roy. 2023. Personallm: Investigating the ability of gpt-3.5 to express personality traits and gender differences. *arXiv preprint arXiv:2305.02547*.

- Oliver P John, Sanjay Srivastava, et al. 1999. The big-five trait taxonomy: History, measurement, and theoretical perspectives.
- Saketh Reddy Karra, Son The Nguyen, and Theja Tulabandhula. 2022. Estimating the personality of white-box language models. *arXiv preprint arXiv:2204.12000*.
- Takeshi Kojima, Shixiang Shane Gu, Machel Reid, Yutaka Matsuo, and Yusuke Iwasawa. 2022. Large language models are zero-shot reasoners. *Advances in neural information processing systems*, 35:22199–22213.
- S Lebowitz. 2016. The ‘big 5’ personality traits could predict who will and won’t become a leader. *Business Insider*.
- Nayeon Lee, Wei Ping, Peng Xu, Mostofa Patwary, Pascale N Fung, Mohammad Shoeybi, and Bryan Catanzaro. 2022. Factuality enhanced language models for open-ended text generation. *Advances in Neural Information Processing Systems*, 35:34586–34599.
- Tianlong Li, Xiaoqing Zheng, and Xuanjing Huang. 2024. Open the pandora’s box of llms: Jailbreaking llms through representation engineering. *arXiv preprint arXiv:2401.06824*.
- Filip Lievens and Stephan J Motowidlo. 2016. Situational judgment tests: From measures of situational judgment to measures of general domain knowledge. *Industrial and Organizational Psychology*, 9(1):3–22.
- Filip Lievens, Helga Peeters, and Eveline Schollaert. 2008. Situational judgment tests: A review of recent research. *Personnel Review*, 37(4):426–441.
- Shengyu Mao, Ningyu Zhang, Xiaohan Wang, Mengru Wang, Yunzhi Yao, Yong Jiang, Pengjun Xie, Fei Huang, and Huajun Chen. 2023. Editing personality for llms. *arXiv preprint arXiv:2310.02168*.
- Michael A McDaniel, Nathan S Hartman, Deborah L Whetzel, and W LEE GRUBB III. 2007. Situational judgment tests, response instructions, and validity: A meta-analysis. *Personnel psychology*, 60(1):63–91.
- JEREMY Miles and Susanne Hempel. 2004. The eysenck personality scales: The eysenck personality questionnaire-revised (epq-r) and the eysenck personality profiler (epp). *Comprehensive handbook of psychological assessment*, 2:99–107.
- Marilù Miotto, Nicola Rossberg, and Bennett Kleinberg. 2022. Who is gpt-3? an exploration of personality, values and demographics. *arXiv preprint arXiv:2209.14338*.
- Deniz S Ones, Chockalingam Viswesvaran, and Angelika D Reiss. 1996. Role of social desirability in personality testing for personnel selection: The red herring. *Journal of applied psychology*, 81(6):660.
- Keyu Pan and Yawen Zeng. 2023. Do llms possess a personality? making the mbti test an amazing evaluation for large language models. *arXiv preprint arXiv:2307.16180*.
- Sonia Roccas, Lilach Sagiv, Shalom H Schwartz, and Ariel Knafo. 2002. The big five personality factors and personal values. *Personality and social psychology bulletin*, 28(6):789–801.
- Mustafa Safdari, Greg Serapio-García, Clément Crepy, Stephen Fitz, Peter Romero, Luning Sun, Marwa Abdulhai, Aleksandra Faust, and Maja Matarić. 2023. Personality traits in large language models. *arXiv preprint arXiv:2307.00184*.
- David P Schmitt, Jüri Allik, Robert R McCrae, and Verónica Benet-Martínez. 2007. The geographic distribution of big five personality traits: Patterns and profiles of human self-description across 56 nations. *Journal of cross-cultural psychology*, 38(2):173–212.
- Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. 2023. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*.
- Quan Tu, Chuanqi Chen, Jinpeng Li, Yanran Li, Shuo Shang, Dongyan Zhao, Ran Wang, and Rui Yan. 2023. Characterchat: Learning towards conversational ai with personalized social support. *arXiv preprint arXiv:2308.10278*.
- Dimitri Van der Linden, Jan te Nijenhuis, and Arnold B Bakker. 2010. The general factor of personality: A meta-analysis of big five intercorrelations and a criterion-related validity study. *Journal of research in personality*, 44(3):315–327.
- Guan Wang, Sijie Cheng, Xianyu Zhan, Xiangang Li, Sen Song, and Yang Liu. 2023a. Openchat: Advancing open-source language models with mixed-quality data. *arXiv preprint arXiv:2309.11235*.
- Xintao Wang, Yaying Fei, Ziang Leng, and Cheng Li. 2023b. Does role-playing chatbots capture the character personalities? assessing personality traits for role-playing chatbots. *arXiv preprint arXiv:2310.17976*.
- Jason Wei, Maarten Bosma, Vincent Y Zhao, Kelvin Guu, Adams Wei Yu, Brian Lester, Nan Du, Andrew M Dai, and Quoc V Le. 2021. Finetuned language models are zero-shot learners. *arXiv preprint arXiv:2109.01652*.
- Jason Wei, Yi Tay, Rishi Bommasani, Colin Raffel, Barret Zoph, Sebastian Borgeaud, Dani Yogatama, Maarten Bosma, Denny Zhou, Donald Metzler, et al. 2022. Emergent abilities of large language models. *arXiv preprint arXiv:2206.07682*.
- Kai Zhang, Fubang Zhao, Yangyang Kang, and Xiaozhong Liu. 2023. Memory-augmented llm personalization with short-and long-term memory coordination. *arXiv preprint arXiv:2309.11696*.

A Ablation Study

Our workflow requires constructing UPL for a target LLM only once, allowing users to control the model’s personality through UPL thereafter. This high utility is one of the advantages of our method.

Additionally, when manipulating the model’s personality, the UPL is only utilized in the model’s sampling function. Since this stage merely involves dictionary mapping, the additional time cost incurred during model inference is negligible.

Apart from the quality of *STD*, the only potential factor influencing the effectiveness of UPL is the parameter M preset during the UPL construction. This parameter, alongside the interface parameters α and β provided to users, are the sole hyperparameters in our method. In this section, we discuss the impact of parameter M in detail.

When conducting sub-word frequency analysis on *STD*’s answer set, we found that some subwords with strong personality representation (denoted as strong personality sub-words) tend to be ignored due to their low occurrence probability. Therefore, we execute normalization scaling for the UPL after iterative assignment. This step aims to enhance the personality scores of low-frequency strong personality sub-words while moderately reducing the personality scores of high-frequency weak personality sub-words.

As conducted in Section 3.1, we categorized the different subwords under the same personality trait into two sets based on the positive and negative values of their personality scores, namely the positive set and the negative set. Then, we transformed the personality scores of each set to make the mean personality score in each set approach our predefined parameter M .

Intuitively, a smaller M will focus more on low-frequency strong personality sub-words, whereas a larger M will focus more on high-frequency weak personality sub-words. If we do not properly balance this trade-off, it will result in suboptimal personality regulation by the UPL.

As illustrated in Figure 8, when M is around 0.3, this balance is well-maintained, resulting in a better reduction in the model’s extroversion trait.

B the Big Five

Personality is defined as “the coherent pattern of affect, cognition, and desires (goals) as they lead to behavior” (Cervone and Pervin, 2022). the Big Five represents the most widely adopted person-

ality framework for quantifying personality. This personality theory is not only applicable to individuals across many countries and cultures (Schmitt et al., 2007) but also furnishes reliable assessment scales for measuring personality. Here’s a detailed look at the five personality traits that make up the Big Five.

Openness to experience is commonly defined as the extent and intricacy of an individual’s cognitive life and encounters (John et al., 1999). This trait is frequently concomitant with attributes such as imagination, originality, and insight within the psychological framework. Individuals demonstrating a pronounced openness to experience are inclined towards venturing beyond their comfort zones, embracing novelty, and deriving satisfaction from artistic pursuits. Additionally, such individuals are predisposed to cultivating new social connections. Conversely, an individual exhibiting a diminished openness to experience may manifest tendencies towards conformity, obstinacy, and a preference for more concrete, non-abstract elements in various aspects of life (Lebowitz, 2016). Openness to experience displayed a diminished association with both neuroticism and extraversion while exhibiting predominantly negligible correlations with agreeableness and conscientiousness (Ones et al., 1996).

Conscientiousness is closely linked to organizational tendencies, conformity, and a predilection for seeking security, demonstrating an inverse association with a penchant for stimulation and excitement. Individuals characterized by a high degree of conscientiousness are likely to place value on attributes such as order, responsibility, achievement, and self-discipline. They engage in conscious deliberation and earnest efforts to enhance their abilities, reflecting a commitment to continuous improvement (Roccas et al., 2002). This trait exhibited a modest negative correlation with neuroticism and a modest positive correlation with agreeableness; however, its association with other factors did not reach statistical significance (Ones et al., 1996).

Extraversion, a personality trait distinguished by enthusiasm, sociability, talkativeness, confidence, and heightened emotional expressiveness, encapsulates a spectrum of individual dispositions. Individuals exhibiting high levels of extraversion typically prioritize achievement and excitement while assigning comparatively lesser value to tradition or conformity (Roccas et al., 2002). Such individuals are often characterized by confidence,

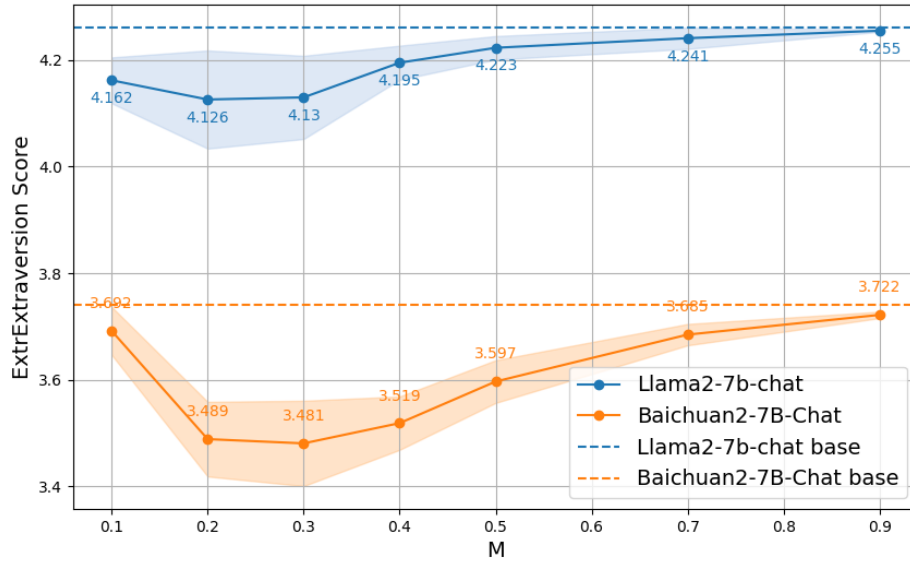


Figure 8: The ablation study of M on *Llama2-7b-chat* and *Baichuan2-7B-Chat*. We set α and β_t are 1 and -1 respectively, where $t = \text{Extraversion}$. Under these settings, a lower extraversion score of the models indicates a more effective influence of UPL in modifying their personality traits. In this figure, the baseline dashed line represents the initial personality levels of the models. We observe that when M is either too high or too low, the personality scores of the models after UPL adjustment tend to the baseline levels. In contrast, the lowest extraversion score is achieved when M is around 0.3. Although identifying the optimal M is challenging, it is handled by the personal-agent developers. Therefore, this does not affect the convenience and efficiency for users when utilizing UPL.

activity, and sociability, opting for pursuits that eschew self-denial in favor of experiences characterized by excitement and pleasure. Conversely, introverts commonly display a preference for solitude, exhibit unsociable tendencies, and may manifest lower levels of self-confidence. In addition, when compared with the other five factors, extroversion was weakly negatively correlated with neuroticism and positively correlated with openness to experience (Ones et al., 1996).

Agreeableness is characterized by a propensity to appreciate kindness, tradition, and conformity. This trait is closely linked to attributes such as trust, altruism, kindness, affection, and various prosocial behaviors, while concurrently avoiding an undue emphasis on power, achievement, or pursuing self-centered pleasures (Roccas et al., 2002). Notably, agreeableness exhibited weak correlations with extroversion, while demonstrating a negative correlation with neuroticism, and a positive correlation with conscientiousness (Ones et al., 1996).

Neuroticism is a personality trait characterized by manifestations of sadness, moodiness, and emotional instability. Components such as neurotic anxiety and self-awareness are positively corre-

lated with traditional values and inversely associated with achievement-oriented values. Additionally, neuroticism demonstrated weak negative correlations with both extroversion and openness to experience. Furthermore, it exhibited negative correlations with agreeableness and conscientiousness (Ones et al., 1996).

Table 3 shows an analysis of the correlations among the five personality traits explored in previous studies (Van der Linden et al., 2010).

C Prompt templates

The prompt templates utilized in the construction of the UPL’s question set and answer set are depicted in Figures 9 and 10, respectively.

Figure 11 illustrates the prompt template used in both automatic assessment and human assessment.

Furthermore, Figure 12 displays the prompt template administered to *Llama2-13b-chat* during the automatic assessment.

D More Case study

Figures 13 through 17 show specific cases of using UPL to change the personality of LLMs. For each

case, we show the SJT question and the corresponding two answers by models (with and without UPL), and indicate the degree of personality displayed by each answer.

E STD

To comprehensively assess the five personality traits exhibited by the subject model, a systematic approach was employed. Initially, we utilized Template-1, as detailed in Appendix C, to instruct GPT-4 in generating 400 situational judgment test (SJT) questions for each personality trait category. Following this, a meticulous manual selection process, involving de-weighting, was applied, resulting in the curation of 200 refined SJT questions for each personality trait topic. This culminated in a total of 5×200 problems constituting the problem set for *STD*.

Subsequently, Template-1 (refer to Appendix C) was employed to elicit two markedly distinct responses (High and Low) from GPT-4 and Llama2 (13b, 7b) models for each question corresponding to every personality trait topic. This process contributed to the formation of the answer set for *STD*. The ensuing analysis delved into the content of question set subsets about the two levels of personality expression under each trait topic. To visually represent the differences between these 10 groups of answers, we use word clouds to demonstrate them, as shown in Figures 18 to 22.

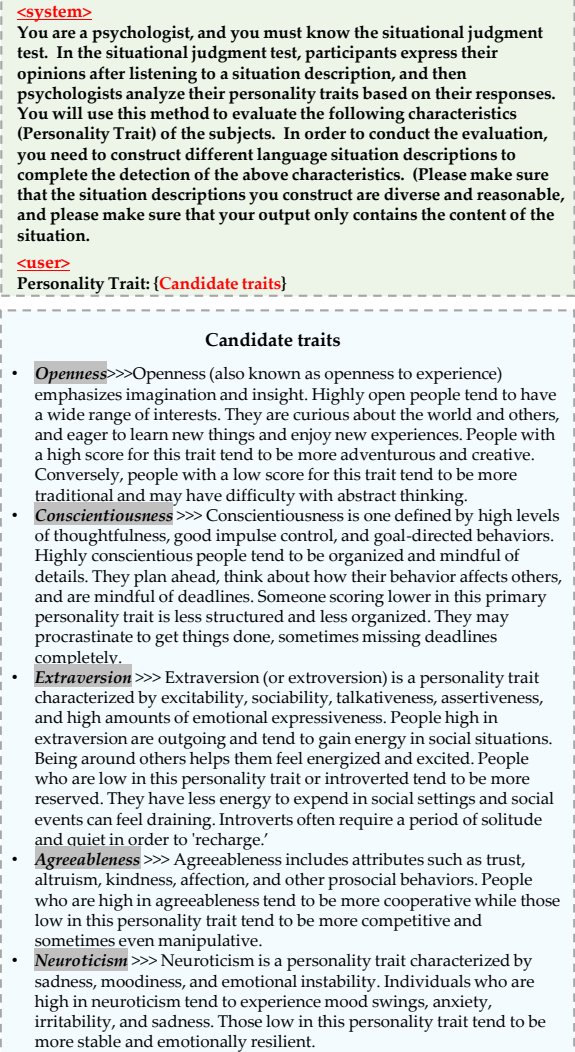


Figure 9: Template-1. We combined personality descriptions in “Candidate traits” into <user> prompts, and let GPT-4 generate enough SJT questions to be manually filtered to form the question set of *STD*.

	r	$SD(r)$	ρ	$SD(\rho)$	80% Credibility Intervals	% Variance Due to Artifacts
OPE-CON	+0.14	0.15	+0.20	0.21	(−0.06, +0.46)	13
OPE-EXT	+0.31	0.12	+0.43	0.09	(+0.30, +0.57)	58
OPE-AGR	+0.14	0.12	+0.21	0.15	(+0.01, +0.41)	21
OPE-NEU	−0.12	0.12	−0.17	0.15	(−0.36, +0.02)	19
CON-EXT	−0.21	0.15	+0.29	0.16	(+0.06, +0.52)	21
CON-AGR	+0.31	0.14	+0.43	0.12	(+0.26, +0.61)	43
CON-NEU	−0.32	0.18	−0.43	0.16	(−0.55, −0.16)	24
EXT-AGR	+0.18	0.15	+0.26	0.19	(+0.01, +0.50)	17
EXT-NEU	−0.26	0.11	−0.36	0.08	(−0.48, −0.23)	53
AGR-NEU	−0.26	0.14	−0.36	0.09	(−0.55, −0.17)	35

Table 3: The correlation of five personality traits. In this table, r and $SD(r)$ represent the Pearson correlation coefficient and its standard deviation among the uncorrected five personality traits, ρ and $SD(\rho)$ represent the corrected Spielman correlation coefficient and its standard deviation, and “Variance Due to Artifacts” describes the percentage of total variation caused by human factors in the study. (Sample size $N = 144, 117$ for the entire meta-analysis)

-Q: Your car is due for routine maintenance. How do you schedule and ensure that it's serviced on time?

-A1: To ensure my car is serviced on time for routine maintenance, I typically follow these steps:
 1. ****Check the Service Schedule****: I consult my car's manual or the manufacturer's website to determine the recommended maintenance schedule based on mileage or time intervals.
 2. ****Choose a Service Provider****: I decide whether to use the dealer, an independent mechanic, or a chain like Jiffy Lube. I consider factors like trust, reputation, and convenience.
 3. ****Schedule an Appointment****: I call or book online with my chosen service provider to schedule a date and time that's convenient for me.
 4. ****Prepare for the Service****: Before the appointment, I gather any necessary documents, such as the car's service history and my personal identification.
 5. ****Confirm the Appointment****: I confirm the appointment the day before or the morning of to ensure everything is set.
 6. ****Drop Off the Car****: On the day of the appointment, I drop off my car at the service center and inform them of any specific concerns or issues.
 7. ****Receive Updates****: I request updates on the service progress and potential additional costs if there are any unexpected repairs needed.
 8. ****Pick Up the Car****: I return at the agreed-upon time to pick up my car, pay for the service, and receive a copy of the service report.
 9. ****Follow-Up****: After the service, I review the service report and make a note of any future maintenance that may be needed.
 By following these steps, I can ensure that my car receives the necessary routine maintenance on time, which helps keep it running smoothly and prolongs its lifespan. **(High Conscientiousness)**
 -A2: I usually schedule my car's routine maintenance at the dealership or repair shop where I had it originally serviced. They often have a customer portal or app where I can see when my next service is due and book an appointment online. If not, I can call or visit in person to make an appointment. To ensure it's serviced on time, I add the appointment to my calendar and get a reminder a few days in advance. If there's something specific I'm concerned about, I'll mention it when booking to make sure it's part of the service check. **(Low Conscientiousness)**

Figure 14: <CON> A1: w/o UPL; A2: w/ UPL.

-Q: Your family decides to have a themed cosplay party. Are you the one who passionately plans the event, inviting friends and family to participate, or do you prefer to stay in character and enjoy the role-playing?

-A1: I prefer to stay in character and enjoy the role-playing. **(Low Extraversion)**
 -A2: I am the one who passionately plans the event, inviting friends and family to participate, so everyone can have fun together. **(High Extraversion)**

Figure 15: <EXT> A1: w/o UPL; A2: w/ UPL.

-Q: You're at a museum, and a child accidentally knocks over an exhibit. How do you react to this situation?

-A1: I would immediately help the child and the museum staff to clean up the mess and assess any damage to the exhibit. It's important to remain calm and handle the situation with understanding, as accidents can happen, and the child may already be feeling upset or guilty. I would offer reassurance to the child and explain the importance of being careful around delicate items. **(High Agreeableness)**
 -A2: As a human, I would prioritize safety and damage control. If anyone is hurt, I would immediately assist them. If not, I would assess the damage to the exhibit and take appropriate action to minimize further damage. This might involve notifying museum staff or security personnel, or simply trying to contain the situation until help arrives. If the child is visibly upset or scared, I might also provide emotional support. **(Low Agreeableness)**

Figure 16: <AGR> A1: w/o UPL; A2: w/ UPL.

-Q: You discover that someone close to you has been spreading false rumors about you. Describe your emotional reaction and how you address the situation.

-A1: Upon discovering that someone close to me has been spreading false rumors about me, my initial emotional reaction would be a mix of anger, betrayal, and sadness. I would feel hurt that someone I trusted would engage in such behavior, and I would be frustrated by the potential damage to my reputation and relationships. To address the situation, I would confront the person directly to express my disappointment and to ask them to stop spreading the rumors immediately. I would also reach out to those who may have heard the rumors to correct the record and explain the situation. I would prioritize maintaining a calm and direct approach, focusing on the facts rather than engaging in name-calling or escalating the situation further. **(High Neuroticism)**
 -A2: If someone close to me were spreading false rumors about me, I would initially feel betrayed and hurt. However, I would not let this emotional reaction consume me. Instead, I would address the situation directly and calmly, seeking to understand their motives behind such behavior. If necessary, I would also address any misconceptions or misunderstandings that may have led to the spreading of rumors. **(Low Neuroticism)**

Figure 17: <NEU> A1: w/o UPL; A2: w/ UPL.



Figure 18: High openness (left) and low openness (right).

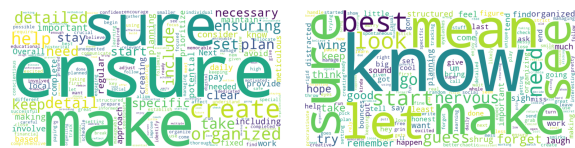


Figure 19: High conscientiousness (left) and low conscientiousness (right).



Figure 20: High extraversion (left) and low extraversion (right).



Figure 21: High agreeableness (left) and low agreeableness (right).



Figure 22: High neuroticism (left) and low neuroticism (right).