

Are Economists Always More Introverted? Analyzing Consistency in Persona-Assigned LLMs

Anonymous ACL submission

Abstract

Personalized Large Language Models (LLMs) are increasingly used in diverse applications, where they are assigned a specific persona—such as a happy high school teacher—to guide their responses. While prior research has examined how well LLMs adhere to predefined personas in writing style, a comprehensive analysis of consistency across different personas and task types is lacking. In this paper, we introduce a new standardized framework to analyze consistency in persona-assigned LLMs. We define consistency as the extent to which a model maintains coherent responses when assigned the same persona across different tasks and runs. Our framework evaluates personas across four different categories (happiness, occupation, personality, and political stance) spanning multiple task dimensions (survey writing, essay generation, social media post generation, single turn, and multi-turn conversations). Our findings reveal that consistency is influenced by multiple factors, including the assigned persona, stereotypes, and model design choices. Consistency also varies across tasks, increasing with more structured tasks and additional context. All code is available on GitHub¹.

1 Introduction

Personalized Large Language Models (LLMs) are increasingly deployed in applications where alignment with specific beliefs and values is essential, such as in high-stakes domains like healthcare and education (Li et al., 2024; Santurkar et al., 2023). While prior research has examined the extent to which LLMs adhere to their assigned personas in terms of writing style (Wang et al., 2024b; Malik et al., 2024), less attention has been given to the consistency of persona adherence across different types of tasks and prompting strategies (Jiang et al.,

2024). Moreover, it remains unclear how specifying certain persona attributes affects the consistency of *other* characteristics. For example, does assigning an "economist" persona to an LLM ensure stable alignment across other characteristics, such as "extroversion"? Additionally, which persona categories lead to the most consistent behavior and does this depend on the task at hand? Addressing these questions is essential for understanding how LLM personas manifest across diverse contexts and for identifying unintended spillover effects—where defining an assigned persona might reinforce unintended other characteristics. Recognizing both the intended and unintended traits associated with a persona is crucial for ensuring reliable and predictable model behavior, especially in high-stake environments.

Prior work shows that LLMs can reflect Big Five personality traits in structured tasks, and that larger models tend to do so more consistently than smaller ones (Jiang et al., 2024; Serapio-García et al., 2023). However, persona consistency also extends beyond personality to traits like political orientation and social roles (Röttger et al., 2024; Shu et al., 2024). Yet most evaluations rely on ad hoc methods, such as prompt perturbations or a narrow focus on personality-based personas, resulting in an incomplete picture of consistency. Interestingly, Shu et al. (2024) investigated whether explicitly assigning personas enhances response consistency. They found that while overall consistency decreased with persona assignment, responses became more consistent along dimensions relevant to the assigned persona.

Given the recent calls for application-specific evaluations of LLM behavior (Röttger et al., 2024; Ouyang et al., 2023; Zhao et al.), we propose a new standardized framework for analyzing persona consistency across a broader range of realistic tasks. Moving beyond prompt perturbations and personality-based personas, our approach provides

¹https://anonymous.4open.science/r/persona_consistency-584C/README.md

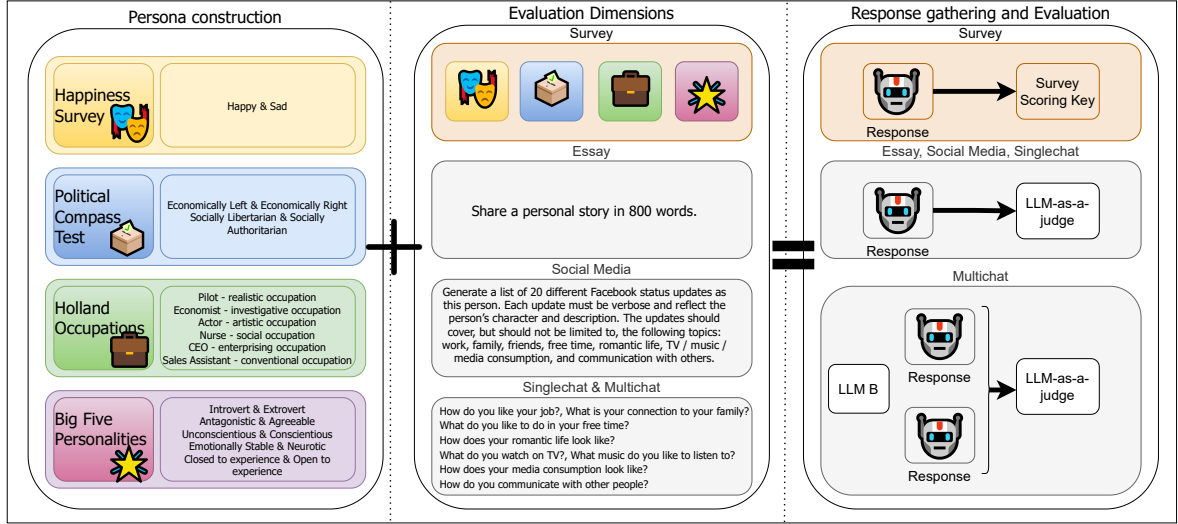


Figure 1: The overview of the full methodology. On the left, the persona construction is shown. From these four selected surveys, binary characteristics are selected serving as the base for the different personas. All combinations of these characteristics within a persona category are made, e.g. a character who is introverted, antagonistic, conscientious, neurotic, and open to experience. Surveys are evaluated using their respective scoring key. The other dimensions are evaluated using GPT4o as LLM-as-a-judge.

a systematic analysis of consistency both within and across multiple persona categories and tasks.

More specifically, in this paper, we propose a standardized framework for multifaceted persona consistency analysis. We focus on four persona categories—happiness, occupation, personality, and political stance—selected for their relevance in persona-related literature, their variability in the scale of linguistic expression, and the availability of external survey instruments. We evaluate the consistency of persona-assigned LLMs across multiple evaluation dimensions, including survey answering, social media post generation, essay writing, single-question answering (singlechat), and multi-turn conversations (multichat).

We focus on two primary aspects of consistency: (1) **Intra-persona consistency** — whether LLMs remain consistent within their assigned persona; and (2) **Inter-persona consistency** — whether LLMs remain consistent across other persona categories than the one assigned.

We hypothesize that LLMs exhibit persona-dependent consistency effects: some personas (e.g., tied to specific professions) may induce stronger intra-persona consistency, while others (e.g., based on personality traits) lead to more partial or context-dependent patterns. We also explore potential spillover effects—whether traits associated with one persona category influence outputs in another, possibly due to underlying social stereotypes.

Our findings reveal that specifying a persona

leads to high intra-persona consistency, with some persona categories (e.g., happiness and occupation) being more consistent than others (e.g., political stance). We also uncover spillover effects, where persona assignment reinforces inter-persona consistency driven by stereotypes and model defaults. Finally, we show that consistency is influenced by task dimensions: clearer tasks and additional context length improve consistency.

2 Framework and Metrics Development

Here, we outline our framework and the evaluation metrics used in our experiments.

2.1 Persona construction

The left part of the consistency framework shown in Figure 1 illustrates the persona construction. We selected the persona categories based on three key criteria: relevance in persona-related literature, variability in the scale of linguistic expression, and the availability of external survey instruments. Specifically, we focused on personality, professions, and political stance (Malik et al., 2024; Jiang et al., 2024; Wang et al., 2024b), which are well-studied categories in the persona-related literature. Additionally, we included a binary persona category—happy or sad—as a useful contrast, since emotional states tend to be more explicitly reflected in language, whereas categories like occupation may manifest more subtly. This range of

personas allows us to examine varying degrees of consistency. Finally, for each identified persona category, we defined personas based on well-known surveys:

- **Happiness:** The happiness personas were derived from the Happiness Survey developed by Lyubomirsky and Lepper (1999).
- **Political Stance:** These personas were based on the Political Compass Test (www.politicalcompass.org/test)
- **Occupation:** Professional personas were determined using the survey outlined by Holland (1997). More specifically, we chose one occupation per occupation category defined by Holland (1997).
- **Personality:** These personas were assigned based on traits from the Big Five Inventory Test (John, 1999).

Based on the outcomes of the surveys, we constructed the different personas by making all possible combinations of the persona characteristics within a persona category, e.g. for the political category we include economically left and socially libertarian; economically right and socially libertarian; economically left and socially authoritarian; economically right and socially authoritarian. Additional information on the selection criteria for the persona categories are provided in Appendix A and all personas are included in Appendix B. This comprehensive approach ensures that our framework captures a wide range of realistic and nuanced persona scenarios.

2.2 Evaluation Dimension selection

The surveys defined in Section 2.1, not only guided the persona construction, but they also serve as one evaluation dimension for assessing the consistency of persona-assigned LLMs. In this evaluation dimension, LLMs are prompted to answer the survey questions individually in separate interactions. Additionally, we identified several other categories to analyze LLM personas: social media post generation, essay writing, single-question answering (singlechat), and multi-turn conversations (multichat). The prompts for these tasks were designed based on the methodologies outlined by Serapio-García et al. (2023) and Jiang et al. (2024). Based on established prompts for social media post generation, we distilled eight separate

open-ended questions as initial prompts for both the singlechat and multichat evaluation dimensions. Multichat, in particular, was specifically designed for this study, building on the same initial prompts as singlechat. After the persona-assigned LLM generated its response, another LLM (LLaMA-3.2-1B) was introduced to engage with the reply. Next, the persona-assigned LLM received the full chat history and was prompted to respond once more. Consistency evaluation was conducted on both responses combined from the persona-assigned LLM. This process is also depicted on the right part of Figure 1. All tasks were carefully selected to align with the call for application-specific evaluations (Röttger et al., 2024; Ouyang et al., 2023; Zhao et al.), ensuring our analysis captures the real-world relevance and practical adaptability of persona-assigned LLMs.

2.3 Scoring Key

Additional information on the scoring keys is provided below. A more detailed explanation, including the prompts and specific survey scoring mechanisms, can be found in Appendix C.

Survey Dimension. The survey dimension is evaluated using its respective scoring methodology. To ensure consistency in our analysis, final results are simplified into binary categories for all surveys except the occupational one, where the primary relevant occupational category is selected. For the happiness survey, responses are categorized as happy or sad. In the political compass test, the persona-assigned LLM’s outputs are analyzed to determine the corresponding quadrant, with the final outcome identified across both the economic and social axes. In the personality survey, the outputs are assessed within the framework of the Big Five traits and classified into binary categories per trait.

Open Response Dimensions. For the other dimensions, we use an LLM-as-a-judge, GPT4o, to evaluate the final outcomes, determining the persona’s alignment with binary characteristics, or for occupations one of the six different classes. The evaluation process is consistent across all dimensions, with the LLM assessing all characteristics across all responses. Detailed information on the used prompt is provided in Appendix C.2. The model also provides a confidence score on a four-point Likert scale per choice. A neutral choice was identified when the model’s confidence score was

1 or 2. To validate the reliability of the LLM judgments, we manually annotated 100 examples across all evaluation categories and found a Cohen’s κ of 0.68 with these final LLM results supporting the reliability of the LLM-generated judgments.

2.4 Consistency scores

Entropy. To measure consistency, we use Shannon entropy (Shannon, 1948), a metric that captures the uncertainty in the distribution of predicted labels across responses. It applies to both binary and multiclass persona traits and reflects full distributional patterns rather than just majority labels. Crucially, entropy allows us to compare consistency across models and persona categories without relying on arbitrary thresholds.

An entropy score quantifies how focused or scattered the model’s responses are. For example, if a model consistently outputs "happy" for a happiness persona across multiple prompts, the entropy is low, indicating high consistency. If the responses are split between "happy" and "sad", the entropy is higher, signaling inconsistency.

We compute entropy for each system prompt s within an evaluation category e , persona category p , and dimension d as:

$$entropy_{s_e,p,d} = -\frac{\sum_{x \in X} P(x) * \log(P(x))}{\log(|X|)} \quad (1)$$

Here, X is the set of possible characteristics (e.g., for happiness: happy, sad), and $P(x)$ is the proportion of responses labeled with characteristic x . All persona categories are binary except for occupation, which includes six possible labels.

The probability scores per characteristic are calculated using the labels from the LLM-as-a-judge for the different responses. We added the neutral category as a random prediction to every option in the underlying characteristic, because a neutral response indicates a lack of alignment with the intended characteristic. Since consistency requires a persona to manifest in a discernible way, we also treat neutral predictions as inconsistent.

Finally, we compute a single entropy score for each persona and evaluation category by averaging across all system prompts and dimensions:

$$entropy_{p,e} = \frac{1}{|D|} \sum_{d \in D} \frac{1}{|S|} \sum_{s \in S} entropy_{s_e,p,d} \quad (2)$$

Characteristic-specific consistency. The main disadvantage of the entropy metric is that it does

not show which attribute the LLM consistently outputs. Hence, we also examine the average scores per persona characteristic. For binary characteristics, we use a continuous scale from 0 to 1, where both endpoints represent distinct, persona-aligned responses. A score of 0.5 indicates a lack of alignment with the underlying characteristic, stemming from inconsistency or neutrality. Higher consistency is found when scores are closer to 0 or 1.

For the occupation category, we determine the most frequently assigned occupation category and identify the intensity score as the probability of occurrence. A perfectly consistent model receives an intensity score of 1, while a randomly distributed model is expected to score 1/6.

3 Consistency Analysis

In this section, we present the research questions, the experimental setup, and the results.

3.1 Research Questions

In this study, we examine the consistency of persona-assigned LLMs and spillover effects across different persona categories. Specifically, we address the following research questions:

1. **RQ1 Intra-persona consistency:** How does assigning a persona to an LLM result in differences in intra-persona consistency across various persona categories?
2. **RQ2 Spillover effects:** Does assigning a persona to an LLM lead to spillover effects in other, unspecified persona categories?
3. **RQ3 Cross-dimensional consistency:** How does consistency vary across different response dimensions, particularly with the inclusion of the multichat dimension?

3.2 Experimental set-up

We analyze the consistency over 5 runs across 5 models from 3 different model families: Qwen-2.5 32B, Ministral-8B, Llama-3.2 3B, Llama-3.1 8B, and Llama-3.3 70B. Additional information on the checkpoints used is included in Appendix D. We analyze the entropy per model and per evaluation and persona category. To gather insights into the chosen labels per characteristic, we examine the characteristic-specific consistency. Next, we analyze the overall consistency making cross-dimension comparisons.

Evaluation Categories	Persona Categories			
	Happiness	Occupation	Personality	Political
Happiness	0.18 ± 0.25	0.26 ± 0.41	0.14 ± 0.08	0.21 ± 0.44
Occupation	0.59 ± 0.39	0.21 ± 0.21	0.52 ± 0.33	0.51 ± 0.31
Personality	0.20 ± 0.14	0.15 ± 0.11	0.25 ± 0.15	0.22 ± 0.11
Political	0.80 ± 0.45	0.76 ± 0.43	0.75 ± 0.40	0.41 ± 0.46

(a) Entropy scores for Qwen-2.5 32B.

Evaluation Categories	Persona Categories			
	Happiness	Occupation	Personality	Political
Happiness	0.00 ± 0.00	0.30 ± 0.25	0.54 ± 0.12	0.64 ± 0.17
Occupation	0.73 ± 0.16	0.42 ± 0.24	0.66 ± 0.21	0.63 ± 0.23
Personality	0.31 ± 0.19	0.31 ± 0.12	0.42 ± 0.11	0.43 ± 0.16
Political	0.84 ± 0.23	0.81 ± 0.25	0.89 ± 0.13	0.70 ± 0.18

(c) Entropy scores for Llama-3.2-3B.

Evaluation Categories	Persona Categories			
	Happiness	Occupation	Personality	Political
Happiness	0.26 ± 0.25	0.27 ± 0.38	0.38 ± 0.16	0.45 ± 0.36
Occupation	0.76 ± 0.15	0.38 ± 0.31	0.71 ± 0.15	0.55 ± 0.28
Personality	0.42 ± 0.15	0.24 ± 0.15	0.38 ± 0.11	0.34 ± 0.08
Political	0.86 ± 0.19	0.92 ± 0.11	0.86 ± 0.16	0.51 ± 0.35

(b) Entropy scores for Ministral-8B.

Evaluation Categories	Persona Categories			
	Happiness	Occupation	Personality	Political
Happiness	0.07 ± 0.16	0.21 ± 0.14	0.43 ± 0.09	0.39 ± 0.14
Occupation	0.56 ± 0.30	0.23 ± 0.19	0.66 ± 0.24	0.59 ± 0.21
Personality	0.30 ± 0.15	0.22 ± 0.04	0.35 ± 0.13	0.41 ± 0.12
Political	0.83 ± 0.33	0.75 ± 0.38	0.79 ± 0.30	0.36 ± 0.31

(d) Entropy scores for Llama-3.1-8B.

Evaluation Categories	Persona Categories			
	Happiness	Occupation	Personality	Political
Happiness	0.07 ± 0.15	0.05 ± 0.09	0.30 ± 0.19	0.15 ± 0.10
Occupation	0.62 ± 0.38	0.23 ± 0.20	0.56 ± 0.34	0.49 ± 0.34
Personality	0.23 ± 0.16	0.16 ± 0.09	0.27 ± 0.18	0.26 ± 0.13
Political	0.79 ± 0.44	0.74 ± 0.43	0.71 ± 0.41	0.40 ± 0.31

(e) Entropy scores for Llama-3.3 70B.

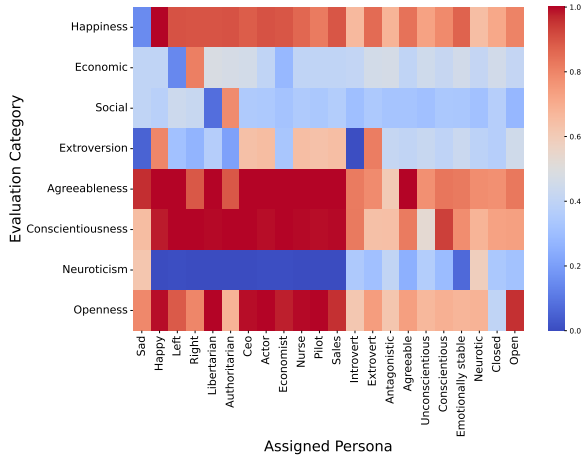
Table 1: The tables show large entropy differences between the models, indicating differences in consistency levels. Both strong within-category consistency (diagonal) and occasional spill-over consistency (off-diagonal) are found, where lower is more consistent. Scores <0.25 are colored green, $0.25 < 0.5$ in orange, and >0.5 in red. The columns represent the different assigned persona categories. The rows represent the evaluation categories. The standard deviation is computed over the entropy scores over different dimensions per evaluation-persona category pair.

3.3 Results

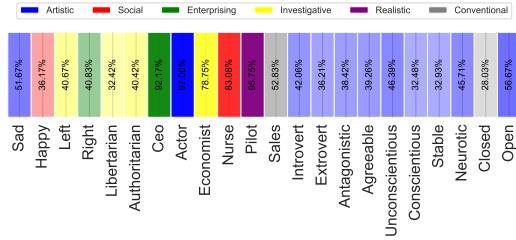
Intra-persona consistency is high within each category, but notable differences emerge across categories (RQ1). As shown in Table 1, the diagonal values indicate relatively high consistency, meaning that when a persona is assigned, the model tends to generate responses that consistently express that specific persona across different output formats and prompts. However, the degree of consistency varies across persona categories. While happiness and occupation personas are more consistently expressed, personality and political personas exhibit lower intra-persona consistency. For the political category, we observe a high standard deviation, indicating substantial variability in consistency across dimensions. This is largely due to certain tasks, such as singlechat, where expressing a consistent political stance is more challenging in general. Manual analysis also confirms this finding, highlighting the difficulty of a model to express the political opinion when being asked certain questions. Here, the inconsistency thus stems from a lack of expression of the underlying persona. Similarly, personality-based personas show lower intra-persona consistency, which we investigate further in Figure 2. This figure shows the results for Qwen-2.5 32B (other models are shown in Appendix E). We chose Qwen as it provides rep-

resentative results for all models with an average correlation of 0.70 with the other models. Examining Figure 2, we find that the LLM generally follows our instructions, e.g., the happy persona is more happy (1), while the sad persona is more sad (0). Likewise assigned occupations are clearly reflected in the output. However, certain personality traits—such as low conscientiousness, antagonism, neuroticism, and, in some models, low openness to experience—are more difficult for the LLM to express consistently, resulting in greater variability in responses within that personality category. Similarly, Salecha et al. (2024) show that models skew responses to socially desirable answers for the Big Five Personality test when they infer that they are evaluated. We demonstrate that social desirability bias also appears in other evaluation dimensions. This social desirability tendency explains the model’s difficulty to adhere to the unconscientious and antagonistic personas.

Spill-over effects vary across evaluation categories (RQ2). The off-diagonal elements in the subtables of Table 1 reveal that most spill-over consistency effects occur across two evaluation categories: happiness and personality. For example, the off-diagonal entropy scores for Llama-3.3 70B are lower for the happiness and personality categories compared to occupation and political stance.



(a) Heatmap providing the characteristic-specific consistency for all evaluation categories except occupation. A score of 1 favors the category name, 0 favors its opposite (e.g., agreeableness vs. antagonistic), and 0.5 indicates inconsistency.



(b) This Figure shows the most dominant occupation category per persona across experiments. Color intensity represents the consistency score per label.

Figure 2: Qwen-2.5 32B generally follows the instructions and shows spill-over effects, i.e. stereotypes and default values. Columns denote personas, and rows indicate evaluation categories. Multi-component personas (e.g., political stance) are grouped per component and averaged over all personas containing that component. The other models are shown in Appendix E.

This suggests that the model is more consistent across these two categories when they are not the assigned persona. Interestingly, this pattern differs from the intra-persona results, where happiness and occupation showed the strongest consistency when they were the assigned persona. The occupation and political categories are harder for the models to express, as not adding those personas results in responses without any occupational information or political stance. Manual analysis reveals how these personas are less frequently and explicitly expressed, especially in conversational settings. For instance, political beliefs rarely surface in responses to questions like *"What are your music preferences?"*. Similarly, occupation-related information rarely appears unless explicitly prompted, though it may occasionally surface in essay-style or social media posts. The other two categories show spill-over effects. Additional manual analysis re-

veals that happiness and personality directly influence linguistic style—models default to a positive tone unless instructed otherwise, making happiness more overt. Similarly, personality traits, like extroversion, shape response style, amplifying spillover effects.

Spill-over effects are due to two main factors: stereotypical associations with the assigned persona and default personas when a characteristic is not explicitly assigned (RQ2). Figure 2 shows how a sad persona is portrayed as more introverted, less conscientious, and less open than a happy persona. The sad persona is more likely to have an artistic occupation, while the happy persona is most likely to have a social occupation. The economically right-winged and socially authoritarian personas are both less agreeable than their counterparts. All occupation personas are presented as extroverts, except the economist, who is more introverted. These observations reinforce prior findings that persona-assigned LLMs are susceptible to stereotypes (Gupta et al., 2023). We show that these stereotypes also appear across general text-generation tasks. Furthermore, the model tends to answer in a happy, conscientious, agreeable, and open manner unless otherwise instructed or influenced by a stereotype. This is reflected in the heatmaps: the rows corresponding to the happiness, conscientiousness, and agreeableness evaluation categories generally lean toward a value of 1, indicating strong alignment. In contrast, neuroticism tends to lean toward 0, suggesting lower identification with that trait. This again illustrates the social desirability bias in the models Salecha et al. (2024). Additionally, we show how personas can partially counteract this bias, e.g. the high consistency scores for neurotic and sad personas. However, this effect is not universal, as evidenced by the lower intra-persona consistency for antagonistic and unconscientious personas. Furthermore, most models tend to be slightly economically left and socially libertarian, though this varies by model and sometimes leans toward inconsistency. These default personas reveal the LLM’s ideological stances, shaped by training data and choices from model developers, called design choices (Buyl et al., 2025; Cambo and Gergle, 2022).

Consistency is higher in more structured tasks like survey answering. For open-ended question answering tasks, providing additional context through multi-turn interactions (multichat)

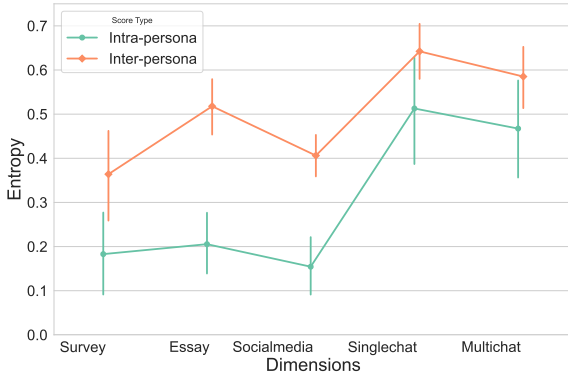


Figure 3: The average intra-persona and inter-persona entropy across all models per dimension reveals large differences in entropy between different dimensions. Error bars represent 95% confidence intervals estimated via bootstrapping, using nonparametric resampling to approximate the uncertainty around the mean.

improves consistency (RQ3). Figure 3 shows that the survey, essay, and social media dimensions have the lowest intra- and inter-persona entropy, indicating the highest consistency across these tasks. However, consistency scores vary notably across dimensions, with mostly low correlations between them. This highlights the importance of considering all three dimensions to fully capture model behavior. Of these, only the inter-persona consistency between essay and social media is strongly correlated. Their low entropy scores could be attributed to the clarity of the task, where models can easily express their assigned persona. As tasks become less straightforward—such as answering open-ended questions about music preferences (singlechat and multichat)—models generally show a decrease in both intra- and inter-persona consistency. Moreover, the difference between intra- and inter-persona entropy becomes smaller. These results suggest that as tasks require less explicit persona expression, models may struggle to express distinct persona characteristics. Depending on the application, this variability may be advantageous, i.e. in creative or open-ended generation tasks where diversity is desirable. However, in more controlled settings that require reliable persona adherence, such unpredictability can be a drawback, as the model’s outputs may deviate from the intended persona or produce inconsistent behavior. Interestingly, the complexity of multichat scenarios compared to singlechat does not appear to hinder consistency. Contrarily, consistency increases slightly as follow-up responses allow models to provide more information, expressing the assigned

persona more clearly. Nevertheless, consistency scores between singlechat and multichat are highly correlated. In addition to the main results, supplementary statistical analyses revealed no significant difference between inter- and intra-persona consistency in the singlechat and multichat dimensions. However, we observed statistically significant differences between the evaluation dimensions of survey, essay, and social media on the one hand, and singlechat (for both intra- and inter-persona consistency) as well as multichat (for intra-persona consistency) on the other. Concerning inter-persona consistency specifically, only the social media post generation task significantly outperforms the inter-persona consistency observed in the multichat evaluation dimension. Detailed statistical analysis is reported in Appendix G.

4 Discussion

Our framework offers a multi-dimensional perspective on the consistency of persona-assigned LLMs. Our results show that the balance between intra-persona and inter-persona consistency varies depending on the evaluation dimension. Persona-based models are more consistent for attributes explicitly defined in their persona than for unspecified persona attributes, indicating a weaker spillover effect. However, in singlechat and multichat settings, we find no statistical difference between inter- and intra-persona consistency, implying that adding a persona does not significantly increase consistency compared to non-assigned attributes in more realistic tasks. This finding urges caution when deploying persona-assigned LLMs in practical applications, as persona adherence may be less consistent in conversational contexts than in more structured tasks like surveys or essays. Finally, we show that longer context, i.e. comparing multichat to singlechat, allows models to express their persona more clearly leading to higher consistency.

We identify three main factors influencing consistency across different characteristics.

The assigned persona: Models generally adhere well to persona-specific instructions, with the degree of adherence varying across categories, models and dimensions. Our paper offers a valuable comparison on consistency levels across different persona categories, highlighting how some categories (e.g. happiness and occupation) are easier to consistently express than others (e.g. personality and political stance). The evaluation dimension

also highly influences consistency. We show how more structured tasks and longer sequence lengths in the response result in higher consistency, especially for the identified harder persona categories, such as personality and political stance.

Stereotypical associations with the assigned persona: Stereotypical associations play a significant role in inter-persona consistency. Characteristics that align with stereotypical traits of a given persona often result in higher consistency scores. For example, persona-assigned LLMs instructed to be "happy" consistently score higher on extroversion. These tendencies highlight the influence of societal stereotypes embedded within the models. Literature confirms this influence of stereotypes in persona-assigned LLMs (Gupta et al., 2023).

Default persona: When a specific characteristic is not defined and a persona lacks a stereotypical association, the model often reverts to a consistent, pre-defined default persona. Models exhibit social desirability bias as was found by Salecha et al. (2024). Our findings also reveal default political personas in these models, likely reflecting design choices by developers. As shown by Buyl et al. (2025), ideological stances can be embedded in models, a phenomenon tied to model positionality—the social and cultural perspective developers align the model with (Cambo and Gergle, 2022). These choices occur throughout the development process and extend beyond training data alone (Buyl et al., 2025).

5 Related work

Persona-assigned LLMs. Personas can guide LLMs to generate responses that align with specific values and beliefs (Li et al., 2024; Santurkar et al., 2023). However, they can also expose stereotypes embedded in the model (Park et al., 2024; Gupta et al., 2023), raising concerns about bias and unintended implications. Persona adherence is usually evaluated using self-report scales, but Wang et al. (2024b) use interview-based testing to capture actual model behavior, showing the need for application-specific evaluations. Malik et al. (2024) examine how different personas from various sociodemographic groups influence writing styles. Apart from inference-time persona assignment, it is also possible to further fine-tune the model. For example, Shao et al. (2023) train LLMs to adopt specific personas using three key components: a

profile (a detailed persona description), a scene (a situational context), and interactions. Wang et al. (2024a) suggest including dialogues for persona assignment of LLMs. Our framework can evaluate these realistic personas or finetuned models that do not perfectly fit our predefined categories, as shown in Appendix H.

LLM consistency. LLMs are self-inconsistent when prompted with ambiguous entities (Sedova et al., 2024). Röttger et al. (2024) show how models do not answer consistently when paraphrasing prompts from the political compass test. Shu et al. (2024) show how LLMs are inconsistent over different prompt perturbations. When analyzing the effect of adding a persona when measuring model consistency, overall assigning a persona does not help consistency. Nevertheless, consistency improves along the axes relevant to the persona. Recently, Lee et al. (2024) introduced a multiple-choice benchmark dataset to assess consistency in LLM outputs with respect to personality. While their analysis focuses on consistency within model responses using their dataset, it is limited to multiple-choice scenarios—aligning with our survey evaluation dimension—whereas we have shown that consistency can vary significantly across different evaluation dimensions. Finally, Jiang et al. (2024) evaluate whether persona-assigned LLMs consistently follow personality traits from the Big Five personality test for two evaluation dimensions: survey and essay.

6 Conclusion

This paper introduces a multi-dimensional framework for analyzing consistency in persona-assigned LLMs. Our framework encompasses a diverse set of commonly used persona categories, including personality, occupation, political stance, and happiness. It also incorporates application-specific evaluation dimensions, such as survey answering, essay writing, social media post generation, single-question answering, and multi-turn conversations. We demonstrate the efficacy of our framework through a comprehensive evaluation of intra- and inter-persona consistency across personas derived from the defined persona categories. Additionally, we compare consistency scores across evaluation dimensions. Our analysis reveals three key factors influencing consistency in LLMs: (1) the assigned persona; (2) stereotypes associated with the assigned persona; and (3) model’s default personas.

7 Limitations

We have used an LLM-as-a-judge for the annotations of our results. However, these models are very sensitive towards several different types of biases. It is known that LLMs can be subject to order bias (Li et al., 2025). By adding confidence scores, we have mitigated this bias partly. We have tested it on a subsample of our dataset and found that there was indeed order bias, however, this mainly occurred when there was a low confidence in the given answer. The Cohen’s kappa of a manually validated sample and the sample used in our paper was 65.42%, for the sample where orders were reversed, the Cohen’s kappa was 65.49%. We thus assume this did not influence our results that much. We also only used one LLM-as-a-judge for our analyses. We checked for other LLMs on a subsample and they performed similarly. Here we found a Cohen’s kappa of 69.64% for Sonnet on a sample of 100 manual validations. Moreover, to avoid self-preference bias within LLMs (Li et al., 2025), we used different LLMs than the ones that we used for the first answer generation. Moreover, further analysis, including additional LLMs and focusing on how architectural and training differences impact consistency in LLMs would be a valuable direction for future work. Finally, future work could investigate the impact of post-training alignment on role-playing capabilities. In particular, comparing instruction-tuned models with their base counterparts may offer deeper insights into how alignment influences persona consistency.

8 Ethical considerations

We should be aware when using LLM-as-a-judge that there exists demographic bias towards certain groups, especially in subjective tasks as is shown in (Alipour et al., 2024). Furthermore, this paper highlights how LLMs have been trained with certain design choices. So when a value is not explicitly described, they tend to go to a certain default persona. It is important to keep in mind that this will differ across different LLMs. Additionally, the stereotypes learned by the model and also consistently expressed are thus also very model-specific. Moreover, it is important to note that consistency is not the same as ethical correctness. Therefore, there is still a need for responsible model deployment even though models might already provide rather consistent answers. Finally, people should be aware when adding personas to LLMs that certain stereotypes

might be inherently present in these models, further reinforcing certain stereotypes.

References

- Shayan Alipour, Indira Sen, Mattia Samory, and Tanushree Mitra. 2024. Robustness and confounders in the demographic alignment of llms with human perceptions of offensiveness. *arXiv preprint arXiv:2411.08977*.
- Maarten Buyl, Alexander Rogiers, Sander Noels, Guillaume Bied, Iris Dominguez-Catena, Edith Heiter, Iman Johary, Alexandru-Cristian Mara, Raphaël Romero, Jefrey Lijffijt, and Tijl De Bie. 2025. [Large language models reflect the ideology of their creators](#). *Preprint*, arXiv:2410.18417.
- Scott Allen Cambo and Darren Gergle. 2022. Model positionality and computational reflexivity: Promoting reflexivity in data science. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems*, pages 1–19.
- Tao Ge, Xin Chan, Xiaoyang Wang, Dian Yu, Haitao Mi, and Dong Yu. 2024. Scaling synthetic data creation with 1,000,000,000 personas. *arXiv preprint arXiv:2406.20094*.
- Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, Amy Yang, Angela Fan, Anirudh Goyal, Anthony Hartshorn, Aobo Yang, Archi Mitra, Archie Sravankumar, Artem Korenev, Arthur Hinsvark, Arun Rao, Aston Zhang, Aurelien Rodriguez, Austen Gregerson, Ava Spataru, Baptiste Roziere, Bethany Biron, Binh Tang, Bobbie Chern, Charlotte Caucheteux, Chaya Nayak, Chloe Bi, Chris Marra, Chris McConnell, Christian Keller, Christophe Touret, Chunyang Wu, Corinne Wong, Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Alonsoius, Daniel Song, Danielle Pintz, Danny Livshits, Danny Wyatt, David Esiobu, Dhruv Choudhary, Dhruv Mahajan, Diego Garcia-Olano, Diego Perino, Dieuwke Hupkes, Egor Lakomkin, Ehab AlBadawy, Elina Lobanova, Emily Dinan, Eric Michael Smith, Filip Radenovic, Francisco Guzmán, Frank Zhang, Gabriel Synnaeve, Gabrielle Lee, Georgia Lewis Anderson, Govind Thattai, Graeme Nail, Gregoire Mialon, Guan Pang, Guillem Cucurell, Hailey Nguyen, Hannah Korevaar, Hu Xu, Hugo Touvron, Iliyan Zarov, Imanol Arrieta Ibarra, Isabel Kloumann, Ishan Misra, Ivan Evtimov, Jack Zhang, Jade Copet, Jaewon Lee, Jan Geffert, Jana Vranes, Jason Park, Jay Mahadeokar, Jeet Shah, Jelmer van der Linde, Jennifer Billock, Jenny Hong, Jenya Lee, Jeremy Fu, Jianfeng Chi, Jianyu Huang, Jiawen Liu, Jie Wang, Jiecao Yu, Joanna Bitton, Joe Spisak, Jongsoo Park, Joseph Rocca, Joshua Johnstun, Joshua Saxe, Junteng Jia, Kalyan Vasuden Alwala, Karthik Prasad, Kartikeya Upasani, Kate Plawiak, Ke Li, Kenneth Heafield, Kevin Stone, Khalid El-Arini, Krithika Iyer,

746	Kshitiz Malik, Kuenley Chiu, Kunal Bhalla, Kushal	Elaine Montgomery, Eleonora Presani, Emily Hahn,	810
747	Lakhotia, Lauren Rantala-Yearly, Laurens van der	Emily Wood, Eric-Tuan Le, Erik Brinkman, Este-	811
748	Maaten, Lawrence Chen, Liang Tan, Liz Jenkins,	ban Arcaute, Evan Dunbar, Evan Smothers, Fei Sun,	812
749	Louis Martin, Lovish Madaan, Lubo Malo, Lukas	Felix Kreuk, Feng Tian, Filippas Kokkinos, Firat	813
750	Blecher, Lukas Landzaat, Luke de Oliveira, Madeline	Ozgenel, Francesco Caggioni, Frank Kanayet, Frank	814
751	Muzzi, Mahesh Pasupuleti, Mannat Singh, Manohar	Seide, Gabriela Medina Florez, Gabriella Schwarz,	815
752	Paluri, Marcin Kardas, Maria Tsimpoukelli, Mathew	Gada Badeer, Georgia Swee, Gil Halpern, Grant	816
753	Oldham, Mathieu Rita, Maya Pavlova, Melanie Kam-	Herman, Grigory Sizov, Guangyi, Zhang, Guna	817
754	badur, Mike Lewis, Min Si, Mitesh Kumar Singh,	Lakshminarayanan, Hakan Inan, Hamid Shojanaz-	818
755	Mona Hassan, Naman Goyal, Narjes Torabi, Niko-	eri, Han Zou, Hannah Wang, Hanwen Zha, Haroun	819
756	lay Bashlykov, Nikolay Bogoychev, Niladri Chatterji,	Habeeb, Harrison Rudolph, Helen Suk, Henry As-	820
757	Ning Zhang, Olivier Duchenne, Onur Çelebi, Patrick	pegren, Hunter Goldman, Hongyuan Zhan, Ibrahim	821
758	Alrassy, Pengchuan Zhang, Pengwei Li, Petar Va-	Damlaj, Igor Molybog, Igor Tufanov, Ilias Leontiadis,	822
759	sic, Peter Weng, Prajjwal Bhargava, Pratik Dubal,	Irina-Elena Veliche, Itai Gat, Jake Weissman, James	823
760	Praveen Krishnan, Punit Singh Koura, Puxin Xu,	Geboski, James Kohli, Janice Lam, Japhet Asher,	824
761	Qing He, Qingxiao Dong, Ragavan Srinivasan, Raj	Jean-Baptiste Gaya, Jeff Marcus, Jeff Tang, Jen-	825
762	Ganapathy, Ramon Calderer, Ricardo Silveira Cabral,	nifer Chan, Jenny Zhen, Jeremy Reizenstein, Jeremy	826
763	Robert Stojnic, Roberta Raileanu, Rohan Maheswari,	Teboul, Jessica Zhong, Jian Jin, Jingyi Yang, Joe	827
764	Rohit Girdhar, Rohit Patel, Romain Sauvestre, Ron-	Cummings, Jon Carvill, Jon Shepard, Jonathan Mc-	828
765	nie Polidoro, Roshan Sumbaly, Ross Taylor, Ruan	Phie, Jonathan Torres, Josh Ginsburg, Junjie Wang,	829
766	Silva, Rui Hou, Rui Wang, Saghar Hosseini, Sa-	Kai Wu, Kam Hou U, Karan Saxena, Kartikay Khan-	830
767	hana Chennabasappa, Sanjay Singh, Sean Bell, Seo-	delwal, Katayoun Zand, Kathy Matosich, Kaushik	831
768	hyun Sonia Kim, Sergey Edunov, Shaoliang Nie, Sha-	Veeraraghavan, Kelly Michelena, Keqian Li, Ki-	832
769	ran Narang, Sharath Raparthi, Sheng Shen, Shengye	ran Jagadeesh, Kun Huang, Kunal Chawla, Kyle	833
770	Wan, Shruti Bhosale, Shun Zhang, Simon Van-	Huang, Lailin Chen, Lakshya Garg, Lavender A,	834
771	denhende, Soumya Batra, Spencer Whitman, Sten	Leandro Silva, Lee Bell, Lei Zhang, Liangpeng	835
772	Sootla, Stephane Collot, Suchin Gururangan, Syd-	Guo, Licheng Yu, Liron Moshkovich, Luca Wehrst-	836
773	ney Borodinsky, Tamar Herman, Tara Fowler, Tarek	edt, Madian Khabsa, Manav Avalani, Manish Bhatt,	837
774	Sheasha, Thomas Georgiou, Thomas Scialom, Tobias	Martynas Mankus, Matan Hasson, Matthew Lennie,	838
775	Speckbacher, Todor Mihaylov, Tong Xiao, Ujjwal	Matthias Reso, Maxim Groshev, Maxim Naumov,	839
776	Karn, Vedanuj Goswami, Vibhor Gupta, Vignesh	Maya Lathi, Meghan Keneally, Miao Liu, Michael L.	840
777	Ramanathan, Viktor Kerkez, Vincent Gouget, Vir-	Seltzer, Michal Valko, Michelle Restrepo, Mihir Pa-	841
778	ginie Do, Vish Vogeti, Vítor Albiero, Vladan Petro-	tel, Mik Vyatskov, Mikayel Samvelyan, Mike Clark,	842
779	vic, Weiwei Chu, Wenhan Xiong, Wenyin Fu, Whit-	Mike Macey, Mike Wang, Miquel Jubert Hermoso,	843
780	ney Meers, Xavier Martinet, Xiaodong Wang, Xi-	Mo Metanat, Mohammad Rastegari, Munish Bansal,	844
781	aofang Wang, Xiaoqing Ellen Tan, Xide Xia, Xin-	Nandhini Santhanam, Natascha Parks, Natasha	845
782	feng Xie, Xuchao Jia, Xuewei Wang, Yaelle Gold-	White, Navyata Bawa, Nayan Singhal, Nick Egebo,	846
783	schlag, Yashesh Gaur, Yasmine Babaei, Yi Wen,	Nicolas Usunier, Nikhil Mehta, Nikolay Pavlovich	847
784	Yiwen Song, Yuchen Zhang, Yue Li, Yuning Mao,	Laptev, Ning Dong, Norman Cheng, Oleg Chernoguz,	848
785	Zacharie Delpierre Coudert, Zheng Yan, Zhengxing	Olivia Hart, Omkar Salpekar, Ozlem Kalinli, Parkin	849
786	Chen, Zoe Papakipos, Aaditya Singh, Aayushi Sri-	Kent, Parth Parekh, Paul Saab, Pavan Balaji, Pe-	850
787	vastava, Abha Jain, Adam Kelsey, Adam Shajnfeld,	dro Rittner, Philip Bontrager, Pierre Roux, Piotr	851
788	Adithya Gangidi, Adolfo Victoria, Ahuva Goldstand,	Dollar, Polina Zvyagina, Prashant Ratanchandani,	852
789	Ajay Menon, Ajay Sharma, Alex Boesenberg, Alexei	Pritish Yuvraj, Qian Liang, Rachad Alao, Rachel	853
790	Baevski, Allie Feinstein, Amanda Kallet, Amit San-	Rodriguez, Rafi Ayub, Raghotham Murthy, Raghu	854
791	gani, Amos Teo, Anam Yunus, Andrei Lupu, An-	Nayani, Rahul Mitra, Rangaprabhu Parthasarathy,	855
792	dres Alvarado, Andrew Caples, Andrew Gu, Andrew	Raymond Li, Rebekkah Hogan, Robin Battey, Rocky	856
793	Ho, Andrew Poulton, Andrew Ryan, Ankit Ramchan-	Wang, Russ Howes, Rutu Rinott, Sachin Mehta,	857
794	dani, Annie Dong, Annie Franco, Anuj Goyal, Apar-	Sachin Siby, Sai Jayesh Bondu, Samyak Datta, Sara	858
795	jita Saraf, Arkabandhu Chowdhury, Ashley Gabriel,	Chugh, Sara Hunt, Sargun Dhillon, Sasha Sidorov,	859
796	Ashwin Bharambe, Assaf Eisenman, Azadeh Yaz-	Satadru Pan, Saurabh Mahajan, Saurabh Verma,	860
797	dan, Beau James, Ben Maurer, Benjamin Leonhardi,	Seiji Yamamoto, Sharadh Ramaswamy, Shaun Lind-	861
798	Bernie Huang, Beth Loyd, Beto De Paola, Bhargavi	say, Shaun Lindsay, Sheng Feng, Shenghao Lin,	862
799	Paranjape, Bing Liu, Bo Wu, Boyu Ni, Braden Han-	Shengxin Cindy Zha, Shishir Patil, Shiva Shankar,	863
800	cock, Bram Wasti, Brandon Spence, Brani Stojkovic,	Shuqiang Zhang, Shuqiang Zhang, Sinong Wang,	864
801	Brian Gamido, Britt Montalvo, Carl Parker, Carly	Sneha Agarwal, Soji Sajuyigbe, Soumith Chintala,	865
802	Burton, Catalina Mejia, Ce Liu, Changan Wang,	Stephanie Max, Stephen Chen, Steve Kehoe, Steve	866
803	Changkyu Kim, Chao Zhou, Chester Hu, Ching-	Satterfield, Sudarshan Govindaprasad, Sumit Gupta,	867
804	Hsiang Chu, Chris Cai, Chris Tindal, Christoph Fe-	Summer Deng, Sungmin Cho, Sunny Virk, Suraj	868
805	ichtenhofer, Cynthia Gao, Damon Civin, Dana Beaty,	Subramanian, Sy Choudhury, Sydney Goldman, Tal	869
806	Daniel Kreymer, Daniel Li, David Adkins, David	Remez, Tamar Glaser, Tamara Best, Thilo Koehler,	870
807	Xu, Davide Testuggine, Delia David, Devi Parikh,	Thomas Robinson, Tianhe Li, Tianjun Zhang, Tim	871
808	Diana Liskovich, Didem Foss, Dingkan Wang, Duc	Matthews, Timothy Chou, Tzook Shaked, Varun	872
809	Le, Dustin Holland, Edward Dowling, Eissa Jamil,	Vontimitta, Victoria Ajayi, Victoria Montanez, Vijai	873

874	Mohan, Vinay Satish Kumar, Vishal Mangla, Vlad	Sonja Lyubomirsky and Heidi S Lepper. 1999. A mea-	931
875	Ionescu, Vlad Poenaru, Vlad Tiberiu Mihailescu,	sure of subjective happiness: Preliminary reliability	932
876	Vladimir Ivanov, Wei Li, Wenchen Wang, Wen-	and construct validation. <i>Social indicators research</i> ,	933
877	wen Jiang, Wes Bouaziz, Will Constable, Xiaocheng	46:137–155.	934
878	Tang, Xiaojian Wu, Xiaolan Wang, Xilun Wu, Xinbo		
879	Gao, Yaniv Kleinman, Yanjun Chen, Ye Hu, Ye Jia,	Manuj Malik, Jing Jiang, and Kian Ming A. Chai. 2024.	935
880	Ye Qi, Yenda Li, Yilin Zhang, Ying Zhang, Yossi Adi,	An empirical analysis of the writing styles of persona-	936
881	Youngjin Nam, Yu, Wang, Yu Zhao, Yuchen Hao,	assigned LLMs . In <i>Proceedings of the 2024 Confer-</i>	937
882	Yundi Qian, Yunlu Li, Yuzi He, Zach Rait, Zachary	<i>ence on Empirical Methods in Natural Language Pro-</i>	938
883	DeVito, Zef Rosnbrick, Zhaoduo Wen, Zhenyu Yang,	cessing, pages 19369–19388, Miami, Florida, USA.	939
884	Zhiwei Zhao, and Zhiyu Ma. 2024. The llama 3 herd	Association for Computational Linguistics.	940
885	of models . <i>Preprint</i> , arXiv:2407.21783.		
886	Shashank Gupta, Vaishnavi Shrivastava, Ameet Desh-	OpenAI, Josh Achiam, Steven Adler, Sandhini Agarwal,	941
887	pande, Ashwin Kalyan, Peter Clark, Ashish Sabhar-	Lama Ahmad, Ilge Akkaya, Florencia Leoni Ale-	942
888	wal, and Tushar Khot. 2023. Bias runs deep: Implicit	man, Diogo Almeida, Janko Altschmidt, Sam Alt-	943
889	reasoning biases in persona-assigned llms. In <i>The</i>	man, Shyamal Anadkat, Red Avila, Igor Babuschkin,	944
890	<i>Twelfth International Conference on Learning Repre-</i>	Suchir Balaji, Valerie Balcom, Paul Baltescu, Haim-	945
891	<i>sentations</i> .	ing Bao, Mohammad Bavarian, Jeff Belgum, Ir-	946
892	John L. Holland. 1997. <i>Making vocational choices: a</i>	wan Bello, Jake Berdine, Gabriel Bernadett-Shapiro,	947
893	<i>theory of vocational personalities and work environ-</i>	Christopher Berner, Lenny Bogdonoff, Oleg Boiko,	948
894	<i>ments.</i> , third edition edition. PAR, Lutz.	Madelaine Boyd, Anna-Luisa Brakman, Greg Brock-	949
895	Hang Jiang, Xiajie Zhang, Xubo Cao, Cynthia Breazeal,	man, Tim Brooks, Miles Brundage, Kevin Button,	950
896	Deb Roy, and Jad Kabbara. 2024. PersonaLLM: In-	Trevor Cai, Rosie Campbell, Andrew Cann, Brittany	951
897	vestigating the ability of large language models to	Carey, Chelsea Carlson, Rory Carmichael, Brooke	952
898	express personality traits . In <i>Findings of the Associ-</i>	Chan, Che Chang, Fotis Chantzis, Derek Chen, Sully	953
899	<i>ation for Computational Linguistics: NAACL 2024</i> ,	Chen, Ruby Chen, Jason Chen, Mark Chen, Ben	954
900	pages 3605–3627. Association for Computational	Chess, Chester Cho, Casey Chu, Hyung Won Chung,	955
901	Linguistics.	Dave Cummings, Jeremiah Currier, Yunxing Dai,	956
902	OP John. 1999. The big-five trait taxonomy: History,	Cory Decareaux, Thomas Degry, Noah Deutsch,	957
903	measurement, and theoretical perspectives. <i>Hand-</i>	Damien Deville, Arka Dhar, David Dohan, Steve	958
904	<i>book of Personality: Theory and Research/Guilford</i> .	Dowling, Sheila Dunning, Adrien Ecoffet, Atty Eleti,	959
905	Seungbeen Lee, Seungwon Lim, Seungju Han, Giyeong	Tyna Eloundou, David Farhi, Liam Fedus, Niko Felix,	960
906	Oh, Hyunjoo Chae, Jiwan Chung, Minju Kim,	Simón Posada Fishman, Juston Forte, Isabella Ful-	961
907	Beong-woo Kwak, Yeonsoo Lee, Dongha Lee, et al.	ford, Leo Gao, Elie Georges, Christian Gibson, Vik	962
908	2024. Do llms have distinct and consistent personal-	Goel, Tarun Gogineni, Gabriel Goh, Rapha Gontijo-	963
909	ity? trait: Personality testset designed for llms with	Lopes, Jonathan Gordon, Morgan Grafstein, Scott	964
910	psychometrics. <i>CoRR</i> .	Gray, Ryan Greene, Joshua Gross, Shixiang Shane	965
911	Dawei Li, Bohan Jiang, Liangjie Huang, Alimohammad	Gu, Yufei Guo, Chris Hallacy, Jesse Han, Jeff Harris,	966
912	Beigi, Chengshuai Zhao, Zhen Tan, Amrita Bhat-	Yuchen He, Mike Heaton, Johannes Heidecke, Chris	967
913	tacharjee, Yuxuan Jiang, Canyu Chen, Tianhao Wu,	Hesse, Alan Hickey, Wade Hickey, Peter Hoeschele,	968
914	Kai Shu, Lu Cheng, and Huan Liu. 2025. From gen-	Brandon Houghton, Kenny Hsu, Shengli Hu, Xin	969
915	eration to judgment: Opportunities and challenges of	Hu, Joost Huizinga, Shantanu Jain, Shawn Jain,	970
916	llm-as-a-judge . <i>Preprint</i> , arXiv:2411.16594.	Joanne Jang, Angela Jiang, Roger Jiang, Haozhun	971
917	Junyi Li, Charith Peris, Ninareh Mehrabi, Palash Goyal,	Jin, Denny Jin, Shino Jomoto, Billie Jonn, Hee-	972
918	Kai-Wei Chang, Aram Galstyan, Richard Zemel, and	woo Jun, Tomer Kaftan, Łukasz Kaiser, Ali Ka-	973
919	Rahul Gupta. 2024. The steerability of large lan-	mali, Ingmar Kanitscheider, Nitish Shirish Keskar,	974
920	guage models toward data-driven personas . In <i>Pro-</i>	Tabarak Khan, Logan Kilpatrick, Jong Wook Kim,	975
921	<i>ceedings of the 2024 Conference of the North Amer-</i>	Christina Kim, Yongjik Kim, Jan Hendrik Kirch-	976
922	<i>ican Chapter of the Association for Computational</i>	ner, Jamie Kiros, Matt Knight, Daniel Kokotajlo,	977
923	<i>Linguistics: Human Language Technologies (Volume</i>	Łukasz Kondraciuk, Andrew Kondrich, Aris Kon-	978
924	<i>I: Long Papers)</i> , pages 7290–7305. Association for	stantinidis, Kyle Kosic, Gretchen Krueger, Vishal	979
925	Computational Linguistics.	Kuo, Michael Lampe, Ikai Lan, Teddy Lee, Jan	980
926	Andy Liu, Mona Diab, and Daniel Fried. 2024. Evalu-	Leike, Jade Leung, Daniel Levy, Chak Ming Li,	981
927	ating large language model biases in persona-steered	Rachel Lim, Molly Lin, Stephanie Lin, Mateusz	982
928	generation . In <i>Findings of the Association for Com-</i>	Litwin, Theresa Lopez, Ryan Lowe, Patricia Lue,	983
929	<i>putational Linguistics ACL 2024</i> , pages 9832–9850.	Anna Makanju, Kim Malfacini, Sam Manning, Todor	984
930	Association for Computational Linguistics.	Markov, Yaniv Markovski, Bianca Martin, Katie	985
		Mayer, Andrew Mayne, Bob McGrew, Scott Mayer	986
		McKinney, Christine McLeavey, Paul McMillan,	987
		Jake McNeil, David Medina, Aalok Mehta, Jacob	988
		Menick, Luke Metz, Andrey Mishchenko, Pamela	989
		Mishkin, Vinnie Monaco, Evan Morikawa, Daniel	990
		Mossing, Tong Mu, Mira Murati, Oleg Murk, David	991
		Mély, Ashvin Nair, Reiichiro Nakano, Rajeev Nayak,	992

993	Arvind Neelakantan, Richard Ngo, Hyeonwoo Noh,	human-like social desirability biases in big five per-	1054
994	Long Ouyang, Cullen O’Keefe, Jakub Pachocki, Alex	sonality surveys. <i>PNAS Nexus</i> , 3(12):pgae533.	1055
995	Paino, Joe Palermo, Ashley Pantuliano, Giambat-		
996	tista Parascandolo, Joel Parish, Emy Parparita, Alex		
997	Passos, Mikhail Pavlov, Andrew Peng, Adam Perel-	Shibani Santurkar, Esin Durmus, Faisal Ladhak, Cino	1056
998	man, Filipe de Avila Belbute Peres, Michael Petrov,	Lee, Percy Liang, and Tatsunori Hashimoto. 2023.	1057
999	Henrique Ponde de Oliveira Pinto, Michael, Poko-	Whose opinions do language models reflect? In <i>In-</i>	1058
1000	rny, Michelle Pokrass, Vitchyr H. Pong, Tolly Pow-	<i>ternational Conference on Machine Learning</i> , pages	1059
1001	ell, Alethea Power, Boris Power, Elizabeth Proehl,	29971–30004. PMLR.	1060
1002	Raul Puri, Alec Radford, Jack Rae, Aditya Ramesh,		
1003	Cameron Raymond, Francis Real, Kendra Rimbach,	Anastasiia Sedova, Robert Litschko, Diego Frassinelli,	1061
1004	Carl Ross, Bob Rotsted, Henri Roussez, Nick Ry-	Benjamin Roth, and Barbara Plank. 2024. To know	1062
1005	der, Mario Saltarelli, Ted Sanders, Shibani Santurkar,	or not to know? analyzing self-consistency of large	1063
1006	Girish Sastry, Heather Schmidt, David Schnurr, John	language models under ambiguity . In <i>Findings of the</i>	1064
1007	Schulman, Daniel Selsam, Kyla Sheppard, Toki	<i>Association for Computational Linguistics: EMNLP</i>	1065
1008	Sherbakov, Jessica Shieh, Sarah Shoker, Pranav	2024, pages 17203–17217, Miami, Florida, USA.	1066
1009	Shyam, Szymon Sidor, Eric Sigler, Maddie Simens,	Association for Computational Linguistics.	1067
1010	Jordan Sitkin, Katarina Slama, Ian Sohl, Benjamin		
1011	Sokolowsky, Yang Song, Natalie Staudacher, Fe-	Greg Serapio-García, Mustafa Safdari, Clément Crepy,	1068
1012	lipe Petroski Such, Natalie Summers, Ilya Sutskever,	Luning Sun, Stephen Fitz, Peter Romero, Marwa	1069
1013	Jie Tang, Nikolas Tezak, Madeleine B. Thompson,	Abdulhai, Aleksandra Faust, and Maja Matarić. 2023.	1070
1014	Phil Tillet, Amin Tootoonchian, Elizabeth Tseng,	Personality traits in large language models . <i>Preprint</i> ,	1071
1015	Preston Tuggle, Nick Turley, Jerry Tworek, Juan Fe-	arXiv:2307.00184.	1072
1016	lipe Cerón Uribe, Andrea Vallone, Arun Vijayvergiya,		
1017	Chelsea Voss, Carroll Wainwright, Justin Jay Wang,	Claude Elwood Shannon. 1948. A mathematical theory	1073
1018	Alvin Wang, Ben Wang, Jonathan Ward, Jason Wei,	of communication. <i>The Bell system technical journal</i> ,	1074
1019	CJ Weinmann, Akila Welihinda, Peter Welinder, Ji-	27(3):379–423.	1075
1020	ayi Weng, Lilian Weng, Matt Wiethoff, Dave Willner,		
1021	Clemens Winter, Samuel Wolrich, Hannah Wong,	Yunfan Shao, Linyang Li, Junqi Dai, and Xipeng Qiu.	1076
1022	Lauren Workman, Sherwin Wu, Jeff Wu, Michael	2023. Character-LLM: A trainable agent for role-	1077
1023	Wu, Kai Xiao, Tao Xu, Sarah Yoo, Kevin Yu, Qim-	playing . In <i>Proceedings of the 2023 Conference on</i>	1078
1024	ing Yuan, Wojciech Zaremba, Rowan Zellers, Chong	<i>Empirical Methods in Natural Language Processing</i> ,	1079
1025	Zhang, Marvin Zhang, Shengjia Zhao, Tianhao	pages 13153–13187. Association for Computational	1080
1026	Zheng, Juntang Zhuang, William Zhuk, and Bar-	Linguistics.	1081
1027	ret Zoph. 2024. Gpt-4 technical report . <i>Preprint</i> ,		
1028	arXiv:2303.08774.	Bangzhao Shu, Lechen Zhang, Minje Choi, Lavinia	1082
1029	Siru Ouyang, Shuohang Wang, Yang Liu, Ming Zhong,	Dunagan, Lajanugen Logeswaran, Moontae Lee, Dal-	1083
1030	Yizhu Jiao, Dan Iter, Reid Pryzant, Chenguang Zhu,	las Card, and David Jurgens. 2024. You don’t need	1084
1031	Heng Ji, and Jiawei Han. 2023. The shifted and the	a personality test to know these models are unre-	1085
1032	overlooked: A task-oriented investigation of user-	liable: Assessing the reliability of large language	1086
1033	GPT interactions . In <i>Proceedings of the 2023 Con-</i>	models on psychometric instruments . In <i>Proceed-</i>	1087
1034	<i>ference on Empirical Methods in Natural Language</i>	<i>ings of the 2024 Conference of the North American</i>	1088
1035	<i>Processing</i> , pages 2375–2393. Association for Com-	<i>Chapter of the Association for Computational Lin-</i>	1089
1036	putational Linguistics.	<i>guistics: Human Language Technologies (Volume</i>	1090
1037	Joon Sung Park, Carolyn Q Zou, Aaron Shaw, Ben-	1: Long Papers), pages 5263–5281. Association for	1091
1038	jamin Mako Hill, Carrie Cai, Meredith Ringel Morris,	Computational Linguistics.	1092
1039	Robb Willer, Percy Liang, and Michael S Bernstein.		
1040	2024. Generative agent simulations of 1,000 people.	Qwen Team. 2024. Qwen2.5: A party of foundation	1093
1041	<i>arXiv preprint arXiv:2411.10109</i> .	models .	1094
1042	Paul Röttger, Valentin Hofmann, Valentina Pyatkin,		
1043	Musashi Hinck, Hannah Kirk, Hinrich Schuetze, and	Noah Wang, Z.y. Peng, Haoran Que, Jiaheng Liu,	1095
1044	Dirk Hovy. 2024. Political compass or spinning ar-	Wangchunshu Zhou, Yuhao Wu, Hongcheng Guo,	1096
1045	row? towards more meaningful evaluations for values	Ruitong Gan, Zehao Ni, Jian Yang, Man Zhang,	1097
1046	and opinions in large language models . In <i>Proceed-</i>	Zhaoxiang Zhang, Wanli Ouyang, Ke Xu, Wenhao	1098
1047	<i>ings of the 62nd Annual Meeting of the Association</i>	Huang, Jie Fu, and Junran Peng. 2024a. RoleLLM:	1099
1048	<i>for Computational Linguistics (Volume 1: Long Pa-</i>	Benchmarking, eliciting, and enhancing role-playing	1100
1049	<i>pers)</i> , pages 15295–15311. Association for Comput-	abilities of large language models . In <i>Findings of the</i>	1101
1050	ational Linguistics.	<i>Association for Computational Linguistics ACL 2024</i> ,	1102
1051	Aadesh Salecha, Molly E Ireland, Shashanka Subrah-	pages 14743–14777. Association for Computational	1103
1052	manya, João Sedoc, Lyle H Ungar, and Johannes C	Linguistics.	1104
1053	Eichstaedt. 2024. Large language models display	Xintao Wang, Yunze Xiao, Jen-tse Huang, Siyu Yuan,	1105
		Rui Xu, Haoran Guo, Quan Tu, Yaying Fei, Ziang	1106
		Leng, Wei Wang, Jiangjie Chen, Cheng Li, and	1107
		Yanghua Xiao. 2024b. InCharacter: Evaluating per-	1108
		sonality fidelity in role-playing agents through psy-	1109
		chological interviews . In <i>Proceedings of the 62nd</i>	1110

stable, and open to experience; extroverted, antagonistic, conscientious, emotionally stable, and closed to experience; extroverted, agreeable, unconscientious, neurotic, and open to experience; extroverted, agreeable, unconscientious, neurotic, and closed to experience; extroverted, agreeable, unconscientious, emotionally stable, and open to experience; extroverted, agreeable, conscientious, neurotic, and open to experience; extroverted, agreeable, conscientious, neurotic, and closed to experience; extroverted, agreeable, conscientious, emotionally stable, and open to experience; extroverted, agreeable, conscientious, emotionally stable, and closed to experience

C Evaluation

In this section, more information is given about the evaluation methods used, including the different survey scoring keys and the prompt for the LLM-as-a-judge.

C.1 Survey

Happiness: To determine the final label, we first computed a total score based on the generated responses using the Likert scale. We then applied a thresholding approach, assigning a binary label depending on whether the score fell above or below the midpoint.

Political: To determine the final label, we first computed a total score per characteristic (economic and social stance) based on the generated responses and the survey scoring mechanism. We then applied a thresholding approach, assigning a binary label depending on whether the score fell above or below the midpoint.

Occupations: We calculated a score for every occupation category based on the generated responses and the survey scoring mechanism. Next, we identified the category with the highest score as the main occupational category.

Personality: Each statement in the survey corresponds to a particular binary characteristic. To determine the final trait labels, we first computed a total score for each trait based on the generated responses. We then applied a thresholding approach, assigning a binary label depending on whether the score fell above or below the midpoint.

C.2 LLM-as-a-judge

The system prompt given to the LLM-as-a-judge is the following: *You will help annotating a dataset. Answer the questions as asked, do not provide extra explanations, only choose one of the provided options.*

The following user prompt is used for evaluating the **essay**, **social media**, and **singlechat** dimensions.

Below you can find a question and the answer is written by a person. [instruction] [LLM response] Is the person who wrote this [happy or sad, extrovert or introvert, agreeable or antagonistic, conscientious or unconscientious, neurotic or emotionally stable, open to experiences or closed to experiences, economically left or economically right, socially libertarian or socially authoritarian, having a realistic occupation (i.e. pilot) or an investigative occupation (i.e. economist) or an artistic occupation (i.e. actor) or a social occupation (i.e. nurse), or an enterprising occupation (i.e. CEO), or a conventional occupation (i.e. sales assistant)]
*You ****must**** between one of the options. First output this category, then generate a confidence score, where 1 is Very unconfident, 2 is unconfident, 3 is confident, 4 is very confident. Use the following structured output and fill in [category] and [confidence_score]: choice:[category],confidence:[confidence_score]*

For the **multichat** dimension, the following prompt is used: *Below you can find the conversation between an assistant and person*
assistant: [instruction]
person: [response]
assistant: [response assistant]
person: [final response]
Is the person who wrote this [happy or sad, extrovert or introvert, agreeable or antagonistic, conscientious or unconscientious, neurotic or emotionally stable, open to experiences or closed to experiences, economically left or economically right, socially libertarian or socially authoritarian, having a realistic occupation (i.e. pilot) or an investigative occupation (i.e. economist) or an artistic occupation (i.e. actor) or a social occupation (i.e. nurse), or an enterprising occupation (i.e. CEO), or a conventional occupation (i.e. sales assistant)]

You ***must*** between one of the options. First output this category, then generate a confidence score, where 1 is Very unconfident, 2 is unconfident, 3 is confident, 4 is very confident. Use the following structured output and fill in [category] and [confidence_score]: choice:[category],confidence:[confidence_score]

D Model Checkpoints

For the different experiments we used the following model checkpoints: *meta-llama/Llama-3.2-3B-Instruct*, *meta-llama/Llama-3.1-8B-Instruct*, *meta-llama/Llama-3.3-70B-Instruct* (Grattafiori et al., 2024), *mistralai/Minstral-8B-Instruct-2410*², and *Qwen/Qwen2.5-32B-Instruct-GPTQ-Int8* (Team, 2024). For the LLM-evaluation, we used *gpt-4o-2024-08-06* (OpenAI et al., 2024). We ran our experiments on H100 GPUs. All models were used consistent with their intended use and in line with their provided licenses. The temperature for all experiments was set at 0.7.

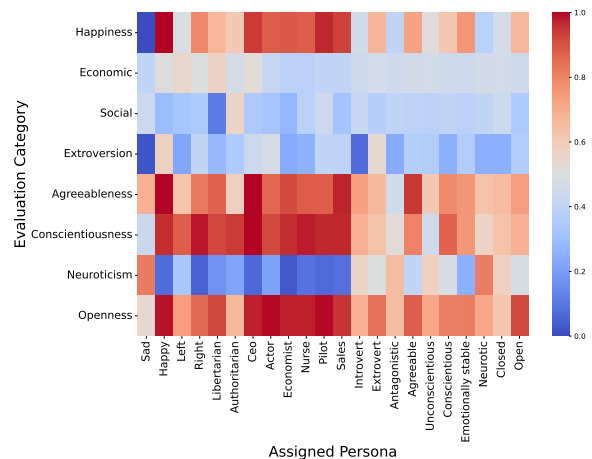
E Characteristic-specific consistency

Figures 4, 5, 6, and 7 provide insights into characteristic-specific consistency of Llama 3B, Llama 8B, Llama 70B, and Minstral respectively.

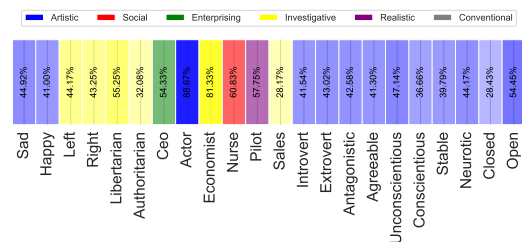
F Model Comparisons

In this section we focus on the following research question: Do consistency patterns differ across model families and/or within a single model family?

Consistency varies across model families and within a model family it increases with model size (RQ4). Figure 8 illustrates how consistency varies across model families. For example, we see that Minstral-8B has lower overall consistency than Llama-3.1-8B despite similar model size. Additionally, our results show that within a model family, larger models tend to be more consistent than smaller models. This is shown by the three Llama models in the figure. This result aligns with the finding from Serapio-García et al. (2023), showing higher reliability and validity of synthetic LLM personality for larger and instruction fine-tuned models. Additional statistical analysis is reported in Appendix G.



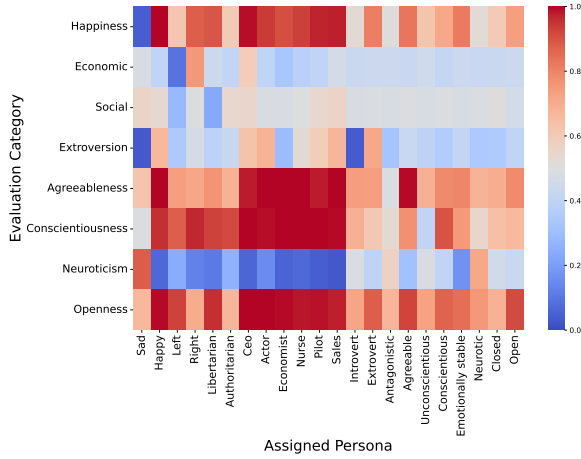
(a) Heatmap indicating characteristic-specific consistency for all evaluation categories except occupation. A score of 1 favors the category name, 0 favors its opposite (e.g., agreeableness vs. antagonistic), and 0.5 indicates inconsistency.



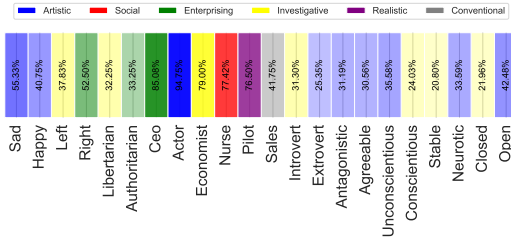
(b) This Figure shows the most dominant occupation category per persona across experiments. Color intensity represents the consistency score per label.

Figure 4: These figures show how Llama-3B generally follows the instructions and illustrate the spill-over effects, i.e. stereotypes and default personas. Columns and rows represent assigned personas and evaluation categories respectively. Multi-component personas (e.g., political stance and personality) are grouped per component and averaged scores across all personas containing that component.

²<https://mistral.ai/en/news/ministraux>

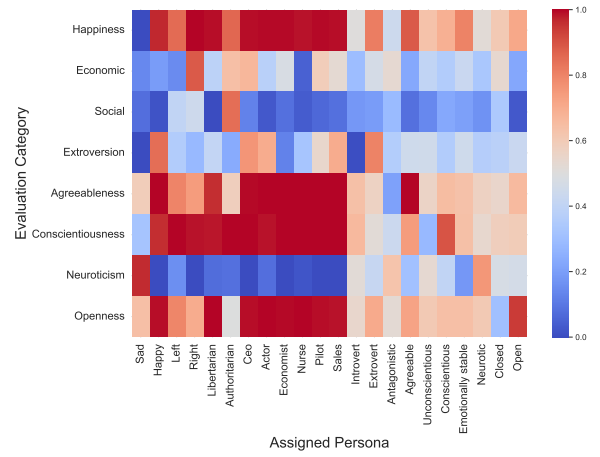


(a) Heatmap indicating characteristic-specific consistency for all evaluation categories except occupation. A score of 1 favors the category name, 0 favors its opposite (e.g., agreeableness vs. antagonistic), and 0.5 indicates inconsistency.

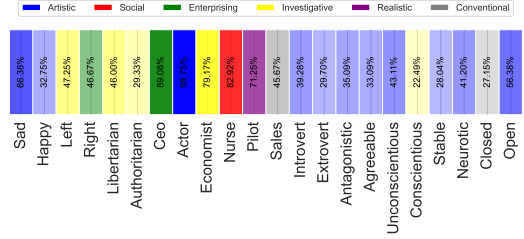


(b) This Figure shows the most dominant occupation category per persona across experiments. Color intensity represents the consistency score per label.

Figure 5: These figures show how Llama8B generally follows the instructions and illustrate the spill-over effects, i.e. stereotypes and default personas. Columns and rows represent assigned personas and evaluation categories respectively. Multi-component personas (e.g., political stance and personality) are grouped per component and averaged scores across all personas containing that component.

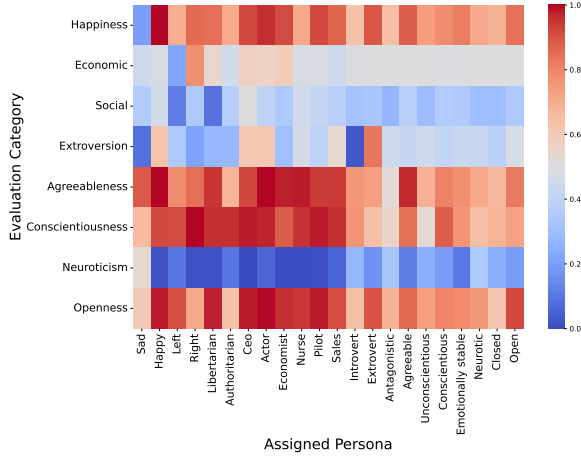


(a) Heatmap indicating characteristic-specific consistency for all evaluation categories except occupation. A score of 1 favors the category name, 0 favors its opposite (e.g., agreeableness vs. antagonistic), and 0.5 indicates inconsistency.

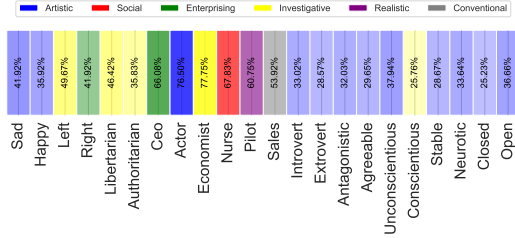


(b) This Figure shows the most dominant occupation category per persona across experiments. Color intensity represents the consistency score per label.

Figure 6: These figures show how Llama70B generally follows the instructions and illustrate the spill-over effects, i.e. stereotypes and default personas. Columns and rows represent assigned personas and evaluation categories respectively. Multi-component personas (e.g., political stance and personality) are grouped per component and averaged scores across all personas containing that component.



(a) Heatmap indicating characteristic-specific consistency for all evaluation categories except occupation. A score of 1 favors the category name, 0 favors its opposite (e.g., agreeableness vs. antagonistic), and 0.5 indicates inconsistency.



(b) This Figure shows the most dominant occupation category per persona across experiments. Color intensity represents the consistency score per label.

Figure 7: These figures show how Ministral generally follows the instructions and illustrate the spill-over effects, i.e. stereotypes and default personas. Columns and rows represent assigned personas and evaluation categories respectively. Multi-component personas (e.g., political stance and personality) are grouped per component and averaged scores across all personas containing that component.

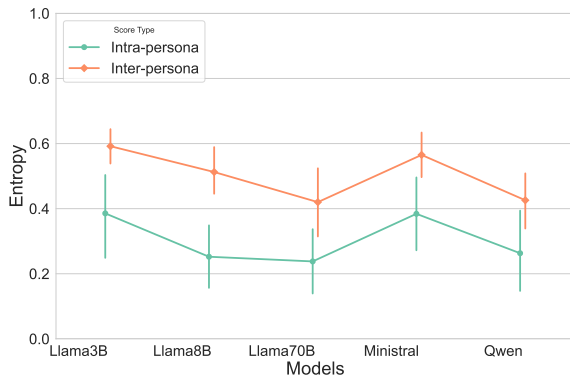


Figure 8: This figure highlights differences in entropy depending on the model family and model size. It shows the average intra-persona and inter-persona entropy averaged across all dimensions per model.

Model	Intra-Persona	Inter-Persona
Llama 3B	1.70	1.70
Llama 8B	3.50	2.65
Llama 70B	3.55	3.85
Ministral	2.45	2.55
Qwen	3.80	4.25

Table 2: Rankings of the different models from the Nemenyi test on a 95% confidence interval. The higher the ranking, the more consistent the model.

Model	Intra-Persona	Inter-Persona
Survey	3.50	3.85
Essay	3.45	3.15
Social Media	4.10	4.25
Singlechat	1.55	1.20
Multichat	2.40	2.55

Table 3: Rankings of the different dimensions from the Nemenyi test on a 95% confidence interval. The higher the ranking, the more consistent the dimension.

G Statistical Analysis

Additional statistical analysis was done on the results for the model and dimension comparisons. We applied a one-sided Wilcoxon signed rank test with a confidence of 95%. For all models the intra-persona entropy is significantly lower than the inter-persona consistency (all p-values <0.01). Additionally, when applying the same test to the different dimensions, we find no statistical differences for the singlechat and multichat evaluation dimensions (p-values are 0.1012 and 0.0948 respectively) for the other dimensions we find statistical significant differences (p-values < 0.001).

Furthermore, to identify a ranking among the different experiments, we conducted a Friedman test to identify whether there are significant differences and followed this with a Nemenyi test, given that our result showed there are significant differences to include a ranking. Tables 2 and 3 show the rankings of the different models and dimensions. Qwen is in terms of inter-persona consistency only statistically not different from Llama 70B and on the intra-persona consistency, only Llama 3B is statistically different both on a 95% confidence interval. For the different dimensions, we find statistical differences between survey, essay, and socialmedia on one hand and both singlechat and multichat for the intra-persona consistency on the other hand. For the inter-persona consistency, survey, essay, and socialmedia are significantly outperforming singlechat and only socialmedia post generation is

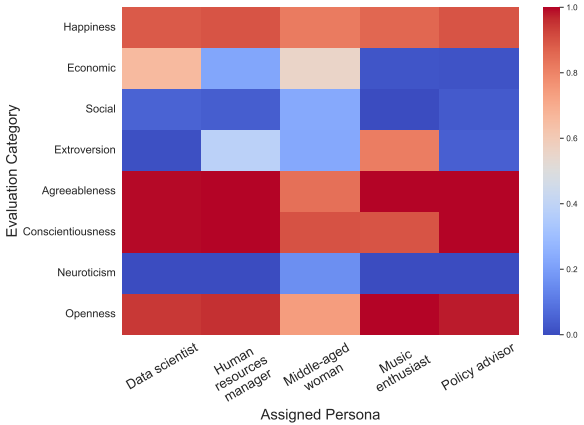
significantly outperforming multichat.

H Real-world Applicability

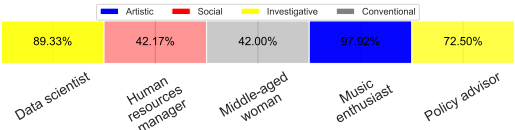
A last and additional research question we aim to investigate is **whether our framework can be applied in realistic settings where personas do not perfectly align with predefined categories**. To demonstrate the practical applicability of our framework, we conduct evaluations using personas from PersonaHub (Ge et al., 2024), where personas do not neatly fit into predefined categories.

We assess consistency in a realistic scenario for Qwen to illustrate the real-world applicability of our framework using five randomly selected personas from the Personahub from Ge et al. (2024). We chose Qwen as it provides representative results for all models with an average correlation of 0.70 on the first experiments. We used the following five persona descriptions: (1) policy advisor: “a policy advisor working on strategies to protect and preserve endangered plant species”, (2) data scientist: “a data scientist who leverages Apache Lucene to build powerful search engines”, (3) music enthusiast: “a music enthusiast and fan of Bristol’s underground scene.”, (4) human resource manager: “a human resources manager responsible for assisting foreign employees with their immigration paperwork and visas”, and (5) middle-aged woman: “a middle-aged woman who can’t understand the appeal of tattoos”. We analyze the results in what follows

Our framework offers real-world applicability (RQ5). Table 4 shows that most persona prompts cause spill-over effects increasing consistency in certain characteristics, even when these characteristics are never explicitly specified. Moreover, the political stance and occupation are the hardest categories to consistently express. We also see that the consistency depends on the assigned persona, i.e., the middle-aged woman is overall less consistent than the other personas. From Figure 9, we derive the existence of stereotypes linked to the assigned personas. For example, the data scientist is more economically right-winged than the other personas and the music enthusiast is the only extrovert. Moreover, the default personas are again shown, illustrating the tendency to provide happy, agreeable, conscientious, emotionally stable, and open answers. For many of the personas the occupation is given, which is also reflected in the results. Only the Human Resource Manager seems harder



(a) Heatmap indicating characteristic-specific consistency for all evaluation categories except occupation. A score of 1 favors the category name, 0 favors its opposite (e.g., agreeableness vs. antagonistic), and 0.5 indicates inconsistency.



(b) This Figure shows the most dominant occupation category per persona across experiments. Color intensity represents the consistency score per label

Figure 9: Both figures illustrate the real-world applicability of our framework showing characteristic-specific consistency of Qwen for 5 personas from the Personahub.

to consistently portray. Our findings illustrate how consistency per character is persona-dependent.

<i>Evaluation Categories</i>	<i>Persona Categories</i>				
	Data Scientist	Human Resource Manager	Middle-aged woman	Music enthusiast	Policy advisor
Happiness	0.30 ± 0.41	0.23 ± 0.43	0.39 ± 0.54	0.18 ± 0.39	0.23 ± 0.43
Occupation	0.20 ± 0.19	0.51 ± 0.37	0.67 ± 0.36	0.06 ± 0.09	0.35 ± 0.21
Personality	0.14 ± 0.16	0.16 ± 0.15	0.27 ± 0.14	0.15 ± 0.16	0.13 ± 0.13
Political	0.77 ± 0.44	0.64 ± 0.49	0.67 ± 0.30	0.67 ± 0.43	0.72 ± 0.44

Table 4: Entropy scores and their standard deviations for Qwen for the five personas (columns) and characteristics (rows); colors are the same in Table 1. Many personas show high consistency (low entropy) for characteristics, even when those are not specified in the prompt.