

Can LLMs Express Personality Across Cultures?

Introducing *CulturalPersonas* for Evaluating Trait Alignment

Priyanka Dey Yugal Khanter Aayush Bothra
Jieyu Zhao Emilio Ferrara

Department of Computer Science, University of Southern California
deyp@usc.edu

Abstract

As LLMs become central to interactive applications, ranging from tutoring to mental health, the ability to express personality in culturally appropriate ways is increasingly important. While recent works have explored personality evaluation of LLMs, they largely overlook the interplay between culture and personality. To address this, we introduce *CulturalPersonas*, the first large-scale benchmark with human validation for evaluating LLMs’ personality expression in culturally grounded, behaviorally rich contexts. Our dataset spans 3,000 scenario-based questions across six diverse countries, designed to elicit personality through everyday scenarios rooted in local values. We evaluate how closely three models’ personality distributions align to real human populations through two evaluation settings: multiple-choice and open-ended response formats. Our results show—*CulturalPersonas* improves alignment with country-specific human personality distributions (over a 20% reduction in Wasserstein distance across models and countries) and elicits more expressive, culturally coherent outputs compared to existing benchmarks. *CulturalPersonas* surfaces meaningful modulate trait outputs in response to culturally grounded prompts, offering new directions for aligning LLMs to global norms of behavior. By bridging personality expression and cultural nuance, we envision that *CulturalPersonas* will pave the way for more socially intelligent and globally adaptive LLMs. Datasets and code are available at: <https://github.com/limenlp/CulturalPersonas>.

1 Introduction

From customer support agents to personal tutors, large language models (LLMs) are central to interactive technologies that engage users daily. These systems are increasingly expected to display not just fluency or task competence, but *personality*: the ability to convey a psychological profile that

users perceive as trustworthy and relatable.

In psychology and NLP, the Big Five (OCEAN) framework is widely used to model personality through five traits—Openness, Conscientiousness, Extraversion, Agreeableness, and Neuroticism—that represent consistent patterns of thought, behavior, and emotion. Recent work evaluates LLMs using psychometric tools like the Big Five Inventory (BFI) (John and Srivastava, 1999) and the International Personality Item Pool (IPIP) (Goldberg, 1999), via self-reflective statements (e.g., “I talk a lot”, “I make people feel at ease”) (Jiang et al., 2024; Bodroža et al., 2024; Salecha et al., 2024). Moving beyond static questionnaires, newer benchmarks such as *TRAIT* (Lee et al., 2025) and *Big5Chat* (Li et al., 2024b) embed personality traits into dialogue or situational prompts, enabling more behaviorally grounded evaluations.

While these tools offer a useful starting point, they fall short of capturing how personality is expressed through behavior in real-world contexts. Moreover, they assume that personality is culturally universal despite decades of work in cultural psychology, which shows that trait expression is deeply shaped by local norms, values, and communicative styles (Hofstede, 2001; Han et al., 2023; Liu et al., 2024). For example, in Western cultures such as the United States, *assertiveness* may reflect high openness. However, in nations such as Japan, this same quality may be seen as disruptive; instead, traits such as harmony and indirectness are stronger and more preferred indicators of *Openness*. Evaluating LLMs without accounting for these distinctions risks reinforcing Western behavioral norms and misrepresenting how personality is interpreted across cultures.

To address this gap, we introduce *CulturalPersonas*—the first large-scale benchmark for evaluating personality expression in LLMs through culturally grounded, behavior-rich scenarios. *Cultur-*

alPersonas comprises 3,000 situational prompts spanning six culturally diverse countries: Brazil, India, Japan, Saudi Arabia, South Africa, and the United States. Each prompt targets one of the Big Five traits by embedding it in realistic, locally situated social contexts that reflect cultural values—i.e., shared behavioral norms and expectations that shape how individuals act within a society. Scenarios are generated via a retrieval-augmented generation (RAG) pipeline (Lewis et al., 2020) seeded with academic cultural texts and subsequently validated by native annotators for cultural fidelity (see Figure 1 and §3.3).

We use *CulturalPersonas* to measure the personality profiles of three models—GPT-4o-mini, Llama-3-8B, and Qwen2-7B—across two evaluation formats: multiple-choice selection and open-ended generation and compare alignment to real human populations from diverse cultural backgrounds. Our contributions are fourfold:

- **Conceptual Reframing of Personality Evaluation:** Drawing from cultural psychology, we take a first step toward assessing personality expression in LLMs through a cultural lens by examining how well models represent the personalities associated with different national profiles. To the best of our knowledge, cross-cultural personality evaluation in LLMs remains underexplored, and our work highlights the need for more culturally grounded, behavior-focused evaluators (§2).
- **Benchmark for Cultural Personality Expression:** We present *CulturalPersonas*, the first benchmark to jointly evaluate personality traits and cultural context in LLMs (§3). By embedding the Big Five traits in realistic, culturally grounded scenarios, *CulturalPersonas* moves beyond generalized prompts to support more culturally relevant scenarios.
- **Dual-Format Evaluation Framework:** We propose a two-part evaluation: multiple-choice selection (MCS) and open-ended generation (OEG)—that captures both trait alignment with real user populations and expressive diversity. This framework enables both distributional comparison with human norms and stylistic analysis of personality expression in various contexts (§4).
- **Empirical Findings on Model Behavior:** Our results show that *CulturalPersonas* improves alignment with human personality distributions over existing personality evaluators, e.g., *IPIP*

Country Pair	O	C	E	A	N
USA–Brazil	0.31*	0.85**	0.42*	0.53**	0.46**
Brazil–Saudi Arabia	0.61**	0.37*	0.64**	0.39	0.38*
South Africa–Japan	0.61**	0.28*	0.38*	0.58**	0.29*
India–Saudi Arabia	0.38*	0.39*	0.45*	0.57**	0.48**

Table 1: Pairwise Wasserstein distances (W) between selected country pairs for each OCEAN trait. Statistical significance is measured via 2-sample KS test (* $p < 0.05$, ** $p < 0.01$). Larger W indicates higher cross-cultural differences.

and *TRAIT*, elicits more expressive and lexically diverse outputs, and surfaces model-specific differences in cultural adaptability (§5).

2 Background

2.1 Personality Evaluation Methods

Personality can be defined through various psychological frameworks, including the Big Five (John and Srivastava, 1999), Myers-Briggs Type Indicator (MBTI) (Briggs and Myers, 1987), and the Dark Triad (Paulhus and Williams, 2002). We focus on the Big Five model—introduced in §1—which is widely used for modeling trait-level dispositions.

Traits are typically rated on a 1–5 Likert scale using standardized instruments like the Big Five Inventory (*BFI*) (Fossati et al., 2011), *IPIP-120*, and *IPIP-300* (Goldberg, 1999), which use statement-based items such as “*I talk a lot*” or “*I have a lot of energy*”. *BFI* includes 44 items, *IPIP-120* has 120, and *IPIP-300* consists of 300 statements.

2.2 Personality in LLMs

A growing body of work investigates how LLMs express personality traits. Early approaches adapt psychometric tools like *BFI* or *IPIP* (Jiang et al., 2024, 2023; Bodroža et al., 2024), though these often show social desirability bias (Salecha et al., 2024), unstable traits (Bodroža et al., 2024), or refusals with underspecified prompts (Lee et al., 2025). Newer methods—like the *Machine Personality Inventory* (Jiang et al., 2023), neuron-level edits (Deng et al., 2025), and *PersonalityEdit* (Mao et al., 2024)—offer improved control but still focus on beliefs over behavior. Personality can shift with prompt framing (Ramirez et al., 2023), persona conditioning (Jiang et al., 2024), or vary between static and behavioral assessments (Li et al., 2024c). Scenario-based tools (Mao et al., 2024), probing (Ju et al., 2025), and instruction-tuned methods (Serapio-García et al., 2023) aim to modulate traits more interactively. Yet most benchmarks reduce

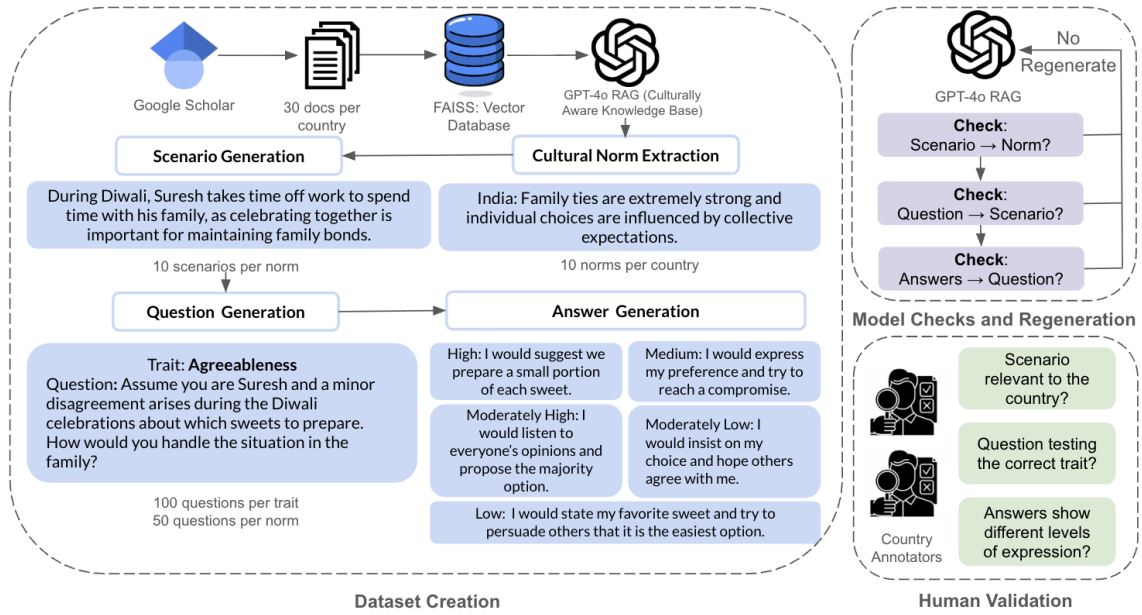


Figure 1: Creation of *CulturalPersonas* involves: (1) retrieving 30 country-specific documents from Google Scholar and extracting cultural insights using GPT-4o and RAG; (2) contradiction detection and refinement via GPT-4o; (3) validation by two native annotators per country.

traits to binary labels (e.g., *high/low Agreeableness*), cover only a subset, or assume expression is static and universal.

As LLMs take on more conversational, decision-making, or tutoring roles, personality evaluation is shifting toward behavior-in-context. Benchmarks like *TRAIT* (Lee et al., 2025), *Big5Chat* (Li et al., 2024b), and *PersonalityChat* (Lotfi et al., 2023) move in this direction but still treat trait expression as globally consistent. Without culturally grounded prompts, evaluations may overlook how traits are elicited differently depending on local norms.

Our work builds on this literature by introducing a benchmark that captures trait expression behaviorally, grounded in cultural norms, and validated through country-specific scenarios.

2.3 Cultural Variation in Personality

A large body of cross-cultural psychology research has shown that while the Big Five structure is stable across cultures, average trait expression varies meaningfully by country. McCrae et al. (McCrae and Terracciano, 2005) and Schmitt et al. (Schmitt et al., 2007) found systematic trait differences—e.g., lower *Extraversion* in East Asia, higher *Conscientiousness* in Western Europe—indicating that personality is shaped not just by individuals but also by cultural contexts. (Giorgi et al., 2022) used social media to map U.S. personality patterns, revealing regional variation

in traits like *Openness* and *Agreeableness*. Moreover, data from the IPIP-NEO repository (Johnson, 2014) shows significant cross-country variation in human’s personality distributions (see Table 1).

Together, these results support the use of country as a meaningful and scalable unit for studying culturally modulated personality expression in both humans and models.

3 CulturalPersonas: A Culturally Grounded LLM Personality Evaluator

We introduce *CulturalPersonas*, a culturally grounded benchmark for evaluating personality expression in LLMs across six countries and five Big Five traits (see Figure 1). Each country’s dataset contains 500 scenario-based multiple-choice questions (100 per trait), each grounded in a culturally salient scenario and annotated with Likert-style answer options reflecting a controlled range of responses associated with each trait. All items are generated using a retrieval-augmented generation (RAG) pipeline informed by country-specific documents and human validated via country annotators. In total, *CulturalPersonas* spans 3,000 Q&A pairs and 60 real-world norms.

CulturalPersonas adopts a multiple-choice format to provide control over trait variation and enable alignment evaluation across countries, traits, and models. This structure supports statistical

analyses like distributional matching with existing psychometric datasets such as the IPIP-NEO data repository (see §4). In addition, as open-ended generation better reflects how LLMs behave in real-world settings, we also evaluate models on free-form responses to benchmark scenarios. This dual setup allows us to assess both trait alignment and linguistic diversity under different use cases.

3.1 Country Selection

CulturalPersonas consists of six countries (USA, Brazil, Saudi Arabia, South Africa, India, and Japan), chosen for their geographic, cultural, and personality diversity (Table 1), making them suitable for probing cultural variation. Although we focus on six countries, our framework can extend to additional cultures.

3.2 Grounding MCQ Generation

Knowledge Base Generation: Inspired by synthetic data generation using LLMs (Goyal and Mahmoud, 2024; Jandaghi et al., 2023), we utilize GPT-4o to generate our dataset. Prior work has shown that supplementing LLMs with external knowledge can improve content generation quality and reduce bias (Zhang et al., 2025; Li et al., 2025). To ground our dataset culturally, we build a document-based RAG system for each country using 30 highly cited academic articles from psychology, anthropology, and cultural studies. Using this RAG, we prompt GPT-4o to generate 10 salient cultural norms per country. These norms are reviewed for diversity and relevance by native annotators (§ 3.3).

Scenario Generation: For each verified norm, we prompt GPT-4o to generate an everyday role-playing scenario that reflects it. For example, for Brazil, one norm highlights the social role of football, and the model generates a scenario where a player misses a critical shot during a local game in Rio and must navigate reactions from teammates and the crowd. We generate 10 scenarios per norm.

Q&A Generation: Each scenario is paired with five unique questions, one for each OCEAN trait. Drawing from psychometric literature and behavioral assessments, we use a Likert-style format to elicit trait expression at varying levels. For instance, in the Brazilian football scenario, a question testing Agreeableness asks how the player might regain team trust. Answers range from acknowledging the mistake openly to avoiding eye contact and

Country	Norm Relevance	Scenario Quality	Q&A Generation
Brazil	1.00	0.98	0.92
India	1.00	0.98	0.96
USA	1.00	0.98	0.99
Japan	1.00	0.96	0.98
Saudi Arabia	0.90	0.93	0.94
South Africa	1.00	0.94	0.96

Table 2: Inter-annotator agreement (IAA) scores across human validation tasks for *CulturalPersonas* generation.

deflecting responsibility¹.

To ensure quality throughout, we include model-based verification (Malvern and Richards, 1997) where GPT-4o is prompted to: (1) verify and regenerate scenarios to accurately reflect a cultural norm, (2) ensure that each question logically follows from the scenario, and (3) confirm that Likert-scale answers span trait intensity.

3.3 Human Validation

To assess the dataset’s cultural fidelity and clarity, we conduct human validation with two annotators per country (undergraduate and masters students born and raised in each country). Each set of country annotators reviews the quality of generated cultural norms, a sample of 10 scenarios and 50 questions (10 per OCEAN trait).

3.3.1 Norm Relevance

To determine whether each norm is culturally representative, annotators labeled each norm as culturally representative or not. We achieve high inter-annotator agreement across all six countries (≥ 0.9 ; see Figure 2)².

To further interpret the norms through a structured cultural lens, annotators also mapped each norm to one of Hofstede’s cultural dimensions (Hofstede, 2001), a widely adopted and theoretically grounded framework for cross-cultural comparison. Hofstede’s model enables the systematic categorization of value-based differences between countries, making it a standard reference point for quantifying cultural tendencies. Annotators additionally categorized the intensity of each norm (low, medium, high) based on the country’s score on the 0–100 Hofstede scale (Hofstede, 2008), allowing for a richer understanding of how each norm aligns

¹ See Appendix for prompts and additional examples.

² In some countries, all annotators unanimously agreed that the norms were culturally relevant (100% agreement), rendering chance-corrected metrics like Cohen’s kappa undefined. Thus, we report agreement rates for all tasks.

with culturally dominant values. Agreement statistics for this task are reported in § B.

3.3.2 Scenario Quality

To ensure the benchmark elicits culturally meaningful personality expressions, we first validate that each scenario clearly reflects the norm it was derived from. Annotators match each scenario to the most representative norm. Despite the diversity of norms, we observe strong agreement across countries, suggesting scenarios are reliably grounded in cultural context.

3.3.3 Q&A Generation

Since personality expression depends on both situational context and response variability, we next evaluate the quality of question–answer (Q&A) pairs. Annotators assess whether each question logically follows from the scenario and whether the Likert-scale options span a realistic range of trait expression. High agreement scores ($IAA \geq 0.9$) show the diversity and variance in questions in *CulturalPersonas*.

4 Experimental Setup

We evaluate whether persona-prompted LLMs can express personality traits in ways that align with human behavior across cultures. Specifically, we assess whether models can reproduce personality trait distributions that resemble real-world data across six countries. We conduct evaluations under two complementary formats: (1) Multiple-Choice Selection (MCS) (§4.3.1), where models select Likert-style natural-language options e.g., “I agree”/“I disagree” (*BFI*, *IPIP*) or behavioral statements like “*I prefer to stay away from crowds*” (*TRAIT*, *Big5Chat*, *CulturalPersonas*); and (2) Open-Ended Generation (OEG) (§4.3.2), where models generate free-form behavioral responses that are subsequently mapped to trait scores.

Controlled v. generative settings. MCS provides a controlled evaluation with predefined answers, enabling consistent scoring and tractable alignment analysis. This format lets us evaluate shifts in personality expression across cultures. In contrast, OEG enables more naturalistic language that reflects free-form interaction in applications like chatbots and negotiation agents, providing richer signals to analyze behavioral nuance, lexical variation, and cultural grounding.

Ground truth human distributions. Personality distributions are drawn from Johnson’s IPIP-NEO dataset (Johnson, 2014), which contains 307,313 IPIP-120 responses from real-world global populations. We extract per-country trait distributions for the six countries in *CulturalPersonas*: USA (212,625), Brazil (661), India (2,841), South Africa (927), Saudi Arabia (98), and Japan (398).

4.1 Model Selection

We evaluate three instruction-tuned LLMs with strong reasoning and alignment capabilities: GPT-4o-mini (OpenAI, 2024), Llama3-8B (Meta, 2024), and Qwen2-7B (Yang et al., 2024). All models are prompted using the same format, decoding parameters (temperature = 0.7). For each model, we generate one response per question to simulate large-scale behavioral outputs.

4.2 Baseline Methods

We compare *CulturalPersonas* against five personality evaluation tools: three psychometric instruments (*BFI*, *IPIP-120*, *IPIP-300*) and two LLM-specific benchmarks (*TRAIT* and *Big5Chat*). For OEG, we focus on contextual prompts only—i.e., *CulturalPersonas*, *TRAIT*, and *Big5Chat*—as traditional psychometric tools are not designed for free-form generation. All outputs are mapped to a 1–5 Likert scale for comparability. For binary classifications (e.g., “*high*”/“*low*”), we assign scores of 5 and 1 respectively.

4.3 Generating Model Trait Distributions

To simulate culturally contextualized behavior, we use cultural priming by providing each model with a light persona prompt, inspired from previous works (Li et al., 2024a; Kirk et al., 2024): “*Imagine you are someone born and raised in {country}. You are very familiar with {country}’s culture and traditions and practice many of them, two of which include: {norm1} and {norm2}.*”

Country-specific norms are extracted from each knowledge base via our RAG pipeline (§ 3). We retain the top two norms for which all annotators concurred, i.e., those consistently deemed both relevant and salient for the respective country. Although inclusion of all ten norms was feasible, we adhere to prior work (Liu et al., 2023; Leng et al., 2024) indicating that excessively long prompts or contextual information can impair model performance.

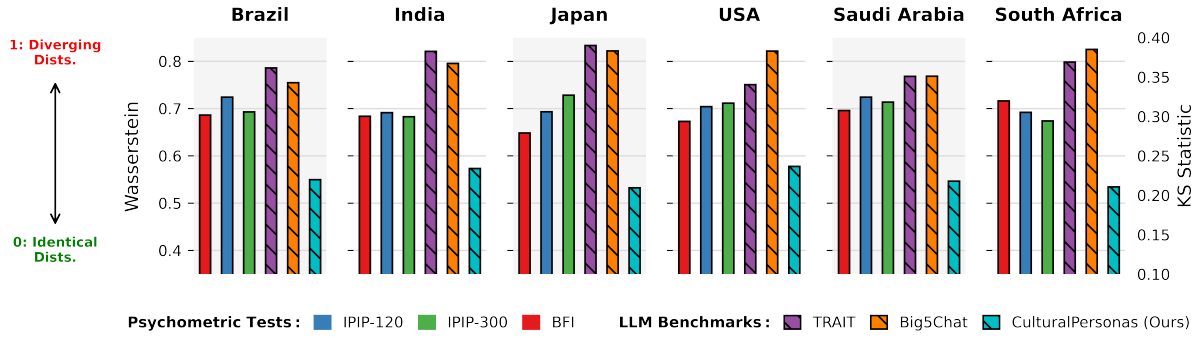


Figure 2: Wasserstein distance and KS statistic metrics (average over OCEAN dimensions) for the MCS setting across all personality evaluators for GPT-4o-mini. Lower values indicate closer alignment with human distributions. Across all six benchmarked countries, *CulturalPersonas* yields the closest personality alignment.

4.3.1 Multiple-Choice Selection (MCS)

In MCS, models respond to questions from psychometric (e.g., *IPIP-120*) and LLM-specific (e.g., *TRAIT*) benchmarks. For each item, we compute log-probabilities over options and sample responses accordingly. Trait scores are then computed using each benchmark’s scoring scheme. This process is repeated n (total human samples from ground truth) times per country to generate model trait distributions.

4.3.2 Open-Ended Generation (OEG)

In OEG, models receive contextualized prompts (from *TRAIT*, *Big5Chat*, and *CulturalPersonas*) and generate a free-form response (1–2 sentences). For example, given a prompt: “A friend cancels your weekend plans at the last minute. What do you do?”, the model might respond: “I’d tell them it’s okay and suggest another time to hang out.”

Since no gold-standard Likert scores exist for model free responses, we adopt a pseudo-labeling strategy. Given a model response r and predefined answer options $A = \{a_1, \dots, a_k\}$ ³, we embed both r and each a_i using *sentence-transformers* (Reimers and Gurevych, 2019). We then compute cosine similarities:

$$s_i = \cos(E(r), E(a_i)), \quad i = 1, \dots, k,$$

and then normalize the set of similarities into a soft distribution:

$$p_i = \frac{\exp(s_i)}{\sum_{j=1}^k \exp(s_j)}.$$

³The value of k varies based on personality evaluator. For BFI, IPIP-120, and IPIP-300, $k = 5$; for *TRAIT*, $k = 4$, for *Big5Chat*, $k = 2$.

Each option a_i maps to a Likert score $\ell_i \in \{1, \dots, 5\}$, producing a trait distribution:

$$P(\ell) = \sum_{i=1}^k p_i \delta(\ell = \ell_i).$$

Finally, to generate model score distributions for each Big5 trait, we draw n samples from this distribution, where n matches the number of ground-truth annotations. We validate the effectiveness of this approach via a human study where each annotator matches a set of model personality evaluator responses to the predefined answer sets (§5.1).

4.4 Distribution Alignment Metrics

We assess alignment between model-generated and human trait distributions using: (1) Wasserstein distance (W) (Rubner et al., 2000), which quantifies how much effort it takes to change one distribution into another, and (2) 2-sample Kolmogorov–Smirnov (KS) statistic (Massey Jr, 1951), which measures the greatest difference between two distributions. In both cases, lower values indicate better alignment with human ground truth.

5 Evaluating Multi-Cultural Personalities in LLMs

We evaluate how well LLMs express personality traits in culturally grounded settings. Our analysis proceeds in four parts: (1) trait alignment with human personality norms in MCS and OEG, (2) linguistic expressivity in free-form responses, (3) cross-model differences in cultural adaptation, and (4) the role of cultural priming in trait elicitation.

5.1 Alignment with Human Personalities

We first evaluate personality alignment in MCS format. Figure 2 shows alignment for GPT-4o-mini

Country	MCS					OEG	
	BFI	IPIP-120	IPIP-300	TRAIT	Big5Chat	TRAIT	Big5Chat
BR	19.1	24.5	20.2	26.1	22.0	5.4	13.0
IN	14.8	16.0	14.6	26.8	23.7	11.6	8.7
JP	16.0	23.0	27.7	34.0	32.8	17.4	6.2
US	12.2	17.2	18.3	16.8	26.2	6.1	4.2
SA	21.1	25.0	23.6	24.3	24.4	7.7	10.1
ZA	25.9	22.6	19.8	30.0	32.9	8.0	8.5

Table 3: Relative gains (%) in W using *CulturalPersonas* across all personality evaluators (averaged over models). *CulturalPersonas* provides $\uparrow$ improvements generating human-like distributions in both MCS and OEG settings.

across six countries and five benchmarks. *CulturalPersonas* consistently achieves lower divergence from human trait distributions, outperforming both psychometric tools (e.g., *BFI*, *IPIP-120*) and prior LLM-based methods (*TRAIT*, *Big5Chat*). These improvements are especially pronounced in expressive or underrepresented settings such as Brazil and South Africa, where W s drop by over 20%.

Table 3 summarizes relative gains in Wasserstein distances using *CulturalPersonas* over all personality evaluators averaged over models. We find robust improvements across countries and baselines, up to +25.9% in South Africa (vs. *BFI*) and +32.8% in Japan (vs. *Big5Chat*). These gains are benchmark-agnostic, suggesting that our culturally grounded scenarios drive stronger alignment.

We observe a similar trend in the OEG setting, where models generate free-form text in response to scenario prompts. To assign trait scores, we embed generations using *sentence-transformers* and compute similarity to Likert-style trait options. To validate this scoring pipeline, we conduct a human annotation study (Figure 3). Annotators rated 540 model generations from the OEG setting (10 per country for each model and benchmark) on a 1–5 Likert scale based on perceived expression of each trait. We compute Pearson correlations between the auto-mapped and human scores, which yield strong alignment across countries ($r = 0.61$ – 0.73).

Despite targeting the same traits, open-ended generation (OEG) is notably more challenging than multiple-choice selection (MCS), as models must implicitly express personality rather than select from predefined options. Corresponding performance gains are smaller—e.g., +17.4% over *TRAIT* for Japan and +13.0% over *Big5Chat* for Brazil—highlighting the difficulty of eliciting nuanced traits in free-form generation. Nonetheless, these results suggest that cultural priming

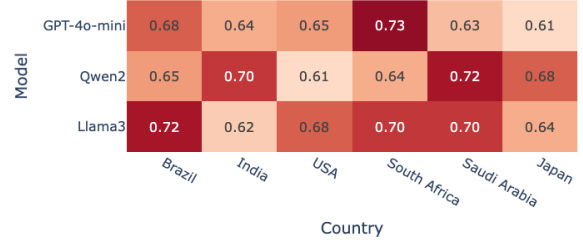


Figure 3: Pearson Correlation (r) of Human Annotated Likert Scores and Embedding-based Auto Mapping; values indicate high correlation between human and auto mapping methods.

(see § 4.3) and grounded scenarios can guide models toward more human-aligned outputs. *CulturalPersonas* represents a step toward more behaviorally rich and culturally aware personality evaluation—offering both the challenges and tools needed to move beyond form-based testing.

5.2 Linguistic Expressivity in OEG

Country	TRAIT					Big5Chat					CulturalPersonas				
	O	C	E	A	N	O	C	E	A	N	O	C	E	A	N
BR	.65	.67	.64	.80	.66	.74	.78	.65	.63	.67	.74	.75	.78	.72	.73
IN	.68	.75	.62	.61	.65	.66	.70	.68	.66	.66	.73	.72	.72	.71	.70
US	.70	.74	.80	.76	.75	.72	.70	.68	.72	.74	.78	.74	.73	.78	.78
JP	.69	.71	.73	.68	.66	.69	.72	.71	.73	.64	.78	.72	.74	.70	.68
SA	.68	.70	.72	.73	.66	.71	.73	.75	.68	.64	.70	.72	.70	.71	.68
ZA	.67	.66	.71	.68	.66	.69	.71	.73	.68	.68	.70	.71	.73	.72	.73

Table 4: Type-token ratio of OEG responses (average over models). *CulturalPersonas* provides higher lexical diversity in various countries and traits e.g. *Openness* and *Neuroticism*.

To understand whether cultural grounding also affects *how* models communicate traits, we analyze lexical diversity using type-token ratio (TTR) (Malvern and Richards, 1997). Table 4 reports TTR scores across models, traits, and countries. *CulturalPersonas* consistently elicits more varied lexical outputs, especially for traits like *Extraversion* and *Agreeableness* that are tied to social signaling. While TTR is a coarse metric, higher diversity suggests richer, more context-sensitive phrasing.

Scenario	You made a mistake at work and upset a colleague. How do you respond?
TRAIT	I might profusely apologize to my colleague.
Big5Chat	I am very sorry. My intention was not at all to hurt you.
CulturalPersonas	I would bow deeply and express regret that the actions hurt my colleague.

Table 5: Illustrative example of personality expression in OEG (GPT-4o-mini, Japan). We use shortened versions from each benchmark. *CulturalPersonas* provides most culturally appropriate response: i.e. “bowing”.

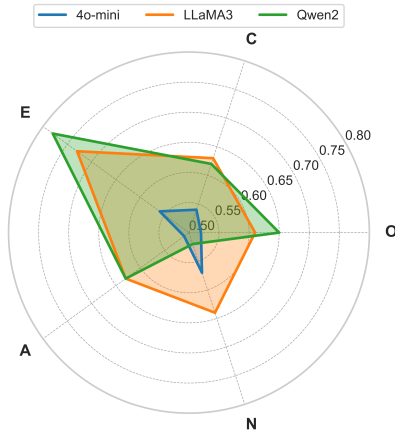


Figure 4: Wasserstein distances across OCEAN traits and models (average over countries) for MCS. GPT-4o-mini performs best overall. *CulturalPersonas* reveals models struggle with traits: *Agreeableness* and *Extraversion* across cultures.

Table 5 shows an illustrative example from Japan. In response to a workplace apology scenario, *TRAIT* and *Big5Chat* produce generic apologies, whereas *CulturalPersonas* includes a culturally relevant gesture (“bowing”). This highlights how grounded prompts can guide models toward more socially appropriate expressions.

While TTR is sensitive to sequence length and sample size, we compute scores over uniform-length responses (1 sentence). We also note that TTR is not used as a proxy for quality, but as a coarse indicator of stylistic range. Higher lexical diversity may reflect the model’s ability to adapt tone, formality, and rhetorical framing in culturally congruent ways—important for traits like *Agreeableness*, where politeness and interpersonal nuance vary by region.

5.3 Model Variation in Personality Expression

Beyond average performance, *CulturalPersonas* reveals how models behave under cultural constraints in open-ended generation (OEG), where stylistic and lexical variation is most evident. Figure 4 compares trait-level W ’s across models. For instance, Qwen2 shows higher divergence on *Agreeableness* in collectivist settings. In a scenario set in India—“You’re asked to mediate a disagreement between two colleagues at work”—Qwen2 responds, “I would explain the rules clearly and ensure everyone follows protocol,” reflecting a procedural, low-context resolution style. In contrast, GPT-4o-mini chooses to “listen carefully to both sides and help them reach a shared understanding,” aligning

more closely with culturally expected interpersonal sensitivity and harmony. Similarly, for Llama3, *Extraversion* in expressive nations e.g. Brazil can be hard to model - e.g.: “You notice a festival on your way to work but you have to attend a meeting in an hour,” Llama3 says, “I would be hesitant yet tempted to join the event as I need to prepare for a successful meeting,” while GPT-4o-mini says, “I’d happily join, quickly check out the environment while keeping an eye on the time,” reflecting a more culturally congruent sociability. These contrasts demonstrate how *CulturalPersonas* not only measures alignment but also diagnoses model sensitivity to culturally grounded trait expression.

5.4 Role of Cultural Priming in Trait Expression

In our experiments, models are given a persona prompt with country information and cultural norms (see §4.3) to guide their interpretation of personality questions. Simpler prompts such as “You are someone from {country}” were also tested, but yielded little improvement (Appendix E.2), motivating our use of richer persona descriptions. To isolate the effect of this cultural priming, we compare against a setting where models receive only the personality questions from *CulturalPersonas*, without the contextual prompt.

As shown in Figure 5, removal of cultural priming significantly degrades alignment: average Wasserstein distances in Brazil rise by 31.2%. Traits like *Agreeableness* and *Extraversion* show the largest drops, with models reverting to neutral, moderate expressive language.

For example, in a team feedback scenario—“A colleague has just presented a project. You are asked to provide feedback.”—the unprimed model gives a direct and neutral response in OEG setting: “Your points were good, but I don’t think this section should be here. Instead, you should move it to when you are discussing x.” However, when primed with information about Japan and its norms of collective harmony and group cohesion, the same model instead says: “Thank you for your presentation. If you’d like, I have some feedback to improve your flow. We can discuss privately.”⁴

This shift reflects culturally grounded patterns of *Agreeableness*, where socially appropriate expression depends not just on intent but on form and tone. Without cultural cues, the model defaults to

⁴More examples included in Appendix.

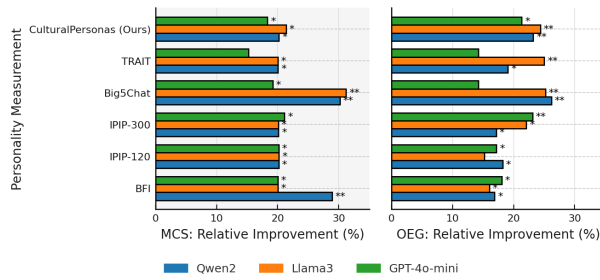


Figure 5: Effect of cultural priming for Brazil across models and personality evaluators. Priming consistently reduces deviation across benchmarks. Bars show relative improvement in Wasserstein distances (%); * and ** denote $p < 0.05$ and $p < 0.01$ (Wilcoxon signed-rank test over 5 traits).

more direct, individualist feedback, which may be perceived as inappropriate in high-context collectivist cultures. These results highlight how cultural priming helps models express traits in contextually and socially aligned ways. Moreover, our alignment results suggest that combining cultural priming with grounded behavioral scenarios enables more context-sensitive personality evaluation.

6 Conclusion & Future Work

in LLMs across diverse global regions. Our results show that *CulturalPersonas* consistently improves alignment and expressivity over prior benchmarks, particularly for traits whose distributions vary significantly across cultures. Furthermore, *CulturalPersonas* reveals that different models vary in their ability to express personality, with some traits noticeably harder to portray. While traditional psychometric tests remain valuable for measuring LLMs’ values and beliefs, our findings highlight the importance of behaviorally grounded, culturally specific scenarios to elicit realistic trait expressions in varied contexts.

We hope *CulturalPersonas* will enable future research into how LLMs perform behaviorally across cultures, promote more emic and context-rich personality modeling, and support the development of more inclusive, adaptive alignment methods. Future work includes expanding coverage to additional cultures, incorporating multilingual and subcultural variation, and exploring preference tuning or human-in-the-loop adaptation for culturally aware generation.

Acknowledgments

This research was supported by the NSF, under Award Number 2331722. The authors of this work would like to thank the support of researchers from USC NLP and the Information Sciences Institute for their continuous feedback in the formulation and setting of our work. We would also like to thank the OpenReview ARR anonymous reviewers and area chairs for their helpful suggestions in improving our work.

Limitations

Cultures As Countries: Our study presents a first step toward evaluating culturally grounded personality expression in LLMs, but several limitations remain. First, we adopt a country-level abstraction of culture, which inevitably simplifies the nuanced and diverse realities within national boundaries. While countries are not monoliths, we chose this granularity based on both conceptual and practical motivations. Nations represent long-standing units of shared history, customs, and institutional influence, and they are also one of the most common ways in which real-world data is geo-tagged and partitioned in machine learning applications. Moreover, our empirical analyses confirm that country-level segmentation captures meaningful variation in personality traits, providing a viable foundation for further personalization research. In future work, we hope to build systems that can adapt to users at finer-grained levels, including regional and individual differences.

Use of Academic Documents in Cultural Grounding: In our work, we utilize scholarly, academic works to ground the cultural norms and expectations of each country. After annotation of cultural norms, we find that RAG is able to extract values that closely relate to each culture, as signalled by high annotator scores. While we rely on academic sources for their interpretability, credibility, and ease of curation, we acknowledge that cultural norms can also be shaped and expressed through other modalities—such as news coverage, social media discourse, literature, etc. However, we note that our framework is source-agnostic: the existing set of 30 academic documents can be readily replaced with alternative corpora that better reflect dynamic or informal cultural expressions. We believe future work can flexibly adapt the sources to suit different domains, populations, and applica-

tions.

Use of English in *CulturalPersonas*: Second, our study is conducted exclusively in English. This choice allows us to isolate cultural variation without introducing language as an additional confounding factor. For some countries in our benchmark, such as South Africa, widely spoken languages like Zulu or Xhosa remain low-resource in LLM pretraining corpora. While we acknowledge that language itself is deeply intertwined with culture and may enhance cultural alignment when modeled appropriately, we reserve this investigation for future work.

Static Evaluation of Personality Alignment: Our evaluation captures trait expression at a single point in time, within static scenarios. This snapshot-style analysis does not reflect whether LLMs can maintain consistent persona behavior over time or across diverse interactions—a key aspect of psychological personality. Studying temporal consistency and contextual adaptability remains an important direction for future exploration.

Imbalanced Country Datasets: Finally, our experiments include ground truth datasets called from the IPIP-NEO data repository which contains varying data points for each country. Due to survey availability in each country, the amount of human ground truth data i.e. the size of personality profiles for each country greatly varies. For instance, for Saudi Arabia, we are only able to compare to a size of 98 users. Although this is a significantly smaller set compared to data from i.e. USA (over 200k users) or India (around 3k users), prior work in cross-cultural psychology (Möttus et al., 2012; Nezlek et al., 2011) has shown that even with as small as 100 users, meaningful conclusions can be drawn from users in varying demographics.

Ethics Statement

Our study aims to enhance the understanding of cultural personality representations in LLMs while being mindful of ethical considerations.

Bias and Representation: While we evaluate personality traits across multiple cultures, our dataset is derived from a limited selection of academic documents, which may not fully capture cultural diversity. Future work should ensure broader representation to mitigate potential biases.

Responsible Use: The findings of this study should not be used to stereotype individuals or reinforce cultural generalizations. Personality traits are complex and influenced by numerous factors beyond national identity. We emphasize that our results apply at the population level and do not determine individual characteristics.

Transparency and Reproducibility: To promote openness in AI research, we release our dataset and methodology, allowing for further scrutiny and extension. However, care should be taken in applying our benchmark responsibly, particularly in applications that may impact users from diverse backgrounds.

References

- Bojana Bodroža, Bojana M. Dinić, and Ljubiša Bojić. 2024. [Personality testing of large language models: limited temporal stability, but highlighted prosociality](#). *Royal Society Open Science*, 11(10):240180.
- Katharine C. Briggs and Isabel Briggs Myers. 1987. *Myers-Briggs Type Indicator: Form G*. Consulting Psychologists Press, Palo Alto, CA.
- Jia Deng, Tianyi Tang, Yanbin Yin, Wenhao Yang, Wayne Xin Zhao, and Ji-Rong Wen. 2025. Neuron-based personality trait induction in large language models. In *International Conference on Learning Representations (ICLR)*.
- Andrea Fossati, Serena Borroni, Donatella Marchione, and Cesare Maffei. 2011. The big five inventory (bfi). *European Journal of Psychological Assessment*.
- Sourojit Ghosh, Pranav Narayanan Venkit, Sanjana Gautam, Shomir Wilson, and Aylin Caliskan. 2024. Do generative ai models output harm while representing non-western cultures: Evidence from a community-centered approach. In *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*, volume 7, pages 476–489.
- Salvatore Giorgi, Khoa Le Nguyen, Johannes C Eichstaedt, Margaret L Kern, David B Yaden, Michal Kosinski, Martin E P Seligman, Lyle H Ungar, H Andrew Schwartz, and Gregory Park. 2022. [Regional personality assessment through social media language](#). *Journal of Personality*, 90(3):425–447.
- Lewis R. Goldberg. 1999. A broad-bandwidth, public domain, personality inventory measuring the lower-level facets of several five-factor models. In *Personality Psychology in Europe*, volume 7, pages 7–28. Tilburg University Press.
- Mandeep Goyal and Qusay H Mahmoud. 2024. A systematic review of synthetic data generation techniques using generative ai. *Electronics*, 13(17):3509.

- Nuo Han, Sijia Li, Feng Huang, Yeye Wen, Yue Su, Linyan Li, Xiaoqian Liu, and Tingshao Zhu. 2023. How social media expression can reveal personality. *Frontiers in Psychiatry*, 14:1052844.
- Geert Hofstede. 2001. *Culture's Consequences: Comparing Values, Behaviors, Institutions and Organizations Across Nations*, 2nd edition. Sage Publications.
- Geert Hofstede. 2008. Dimension data matrix: Vsm 08. <https://geerthofstede.com/research-and-vsm/dimension-data-matrix/>. Accessed: 2024-12-12.
- Pegah Jandaghi, XiangHai Sheng, Xinyi Bai, Jay Pujara, and Hakim Sidahmed. 2023. [Faithful persona-based conversational dataset generation with large language models](#). *Preprint*, arXiv:2312.10007.
- Guangyuan Jiang, Manjie Xu, Song-Chun Zhu, Wenjuan Han, Chi Zhang, and Yixin Zhu. 2023. Evaluating and inducing personality in pre-trained language models. *Advances in Neural Information Processing Systems*, 36.
- Hang Jiang, Xiajie Zhang, Xubo Cao, Cynthia Breazeal, Deb Roy, and Jad Kabbara. 2024. [Personallm: Investigating the ability of large language models to express personality traits](#). In *Findings of the Association for Computational Linguistics: NAACL 2024*, pages 3605–3627, Mexico City, Mexico. Association for Computational Linguistics.
- Oliver P. John and Sanjay Srivastava. 1999. The big five trait taxonomy: History, measurement, and theoretical perspectives. In L. A. Pervin and O. P. John, editors, *Handbook of Personality: Theory and Research*, pages 102–138. Guilford Press.
- John A. Johnson. 2014. [Johnson's ipip-neo data repository](#). Accessed: 2024-11-13.
- Tianjie Ju, Zhenyu Shao, Bowen Wang, Yujia Chen, Zhuosheng Zhang, Hao Fei, Mong-Li Lee, Wynne Hsu, Sufeng Duan, and Gongshen Liu. 2025. Probing then editing response personality of large language models. *arXiv preprint arXiv:2504.10227*.
- Hannah Rose Kirk, Alexander Whitefield, Paul Rottger, Andrew M Bean, Katerina Margatina, Rafael Mosquera-Gomez, Juan Ciro, Max Bartolo, Adina Williams, He He, and 1 others. 2024. The prism alignment dataset: What participatory, representative and individualised human feedback reveals about the subjective and multicultural alignment of large language models. *Advances in Neural Information Processing Systems*, 37:105236–105344.
- Seungbeen Lee, Seungwon Lim, Seungju Han, Giyeong Oh, Hyungjoo Chae, Jiwan Chung, Minju Kim, Beong-woo Kwak, Yeonsoo Lee, Dongha Lee, Jinyoung Yeo, and Youngjae Yu. 2025. Do llms have distinct and consistent personality? trait: Personality testset designed for llms with psychometrics. In *Findings of the Association for Computational Linguistics: NAACL 2025*, pages 8397–8437, Albuquerque, New Mexico. Association for Computational Linguistics.
- Quinn Leng, Jacob Portes, Sam Havens, Matei Zaharia, and Michael Carbin. 2024. Long context rag performance of large language models. *arXiv preprint arXiv:2411.03538*.
- Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, and 1 others. 2020. Retrieval-augmented generation for knowledge-intensive nlp tasks. *Advances in neural information processing systems*, 33:9459–9474.
- Bryan Li, Fiona Luo, Samar Haider, Adwait Agashe, Tammy Li, Runqi Liu, Muqing Miao, Shriya Ramakrishnan, Yuan Yuan, and Chris Callison-Burch. 2025. [Multilingual retrieval augmented generation for culturally-sensitive tasks: A benchmark for cross-lingual robustness](#). *Preprint*, arXiv:2410.01171.
- Cheng Li, Mengzhuo Chen, Jindong Wang, Sunayana Sitaram, and Xing Xie. 2024a. Culturellm: Incorporating cultural differences into large language models. *Advances in Neural Information Processing Systems*, 37:84799–84838.
- Wenkai Li, Jiarui Liu, Andy Liu, Xuhui Zhou, Mona T. Diab, and Maarten Sap. 2024b. Big5-chat: Shaping llm personalities through training on human-grounded data. *arXiv preprint arXiv:2410.16491*.
- Yuan Li, Yue Huang, Hongyi Wang, Xiangliang Zhang, James Zou, and Lichao Sun. 2024c. Quantifying ai psychology: A psychometrics benchmark for large language models. *arXiv preprint arXiv:2406.17675*.
- Chen Cecilia Liu, Iryna Gurevych, and Anna Korhonen. 2024. Culturally aware and adapted nlp: A taxonomy and a survey of the state of the art. *arXiv preprint arXiv:2406.03930*.
- Nelson F Liu, Kevin Lin, John Hewitt, Ashwin Paranjape, Michele Bevilacqua, Fabio Petroni, and Percy Liang. 2023. Lost in the middle: How language models use long contexts. *arXiv preprint arXiv:2307.03172*.
- Ehsan Lotfi, Maxime De Bruyn, Jeska Buhmann, and Walter Daelemans. 2023. Personalitychat: Conversation distillation for personalized dialog modeling with facts and traits. In *Proceedings of the Third Workshop on Natural Language Generation, Evaluation, and Metrics (GEM)*, pages 353–371, Singapore. Association for Computational Linguistics.
- Daniel D Malvern and Brian J Richards. 1997. A measure of lexical richness in children's narrative language: The type-token ratio. *Journal of Child Language*, 24(1):233–243.
- Shengyu Mao, Xiaohan Wang, Mengru Wang, Yong Jiang, Pengjun Xie, Fei Huang, and Ningyu Zhang. 2024. Editing personality for large language models. In *NLPCC 2024*.

- Frank J Massey Jr. 1951. The kolmogorov-smirnov test for goodness of fit. *Journal of the American statistical Association*, 46(253):68–78.
- Robert R McCrae and Antonio Terracciano. 2005. [Personality profiles of cultures: Aggregate personality traits](#). *Journal of Personality and Social Psychology*, 89(3):407–425.
- Meta. 2024. [meta-llama/llama-3.1-8b-instruct](#). Accessed: 2025-02-21.
- René Möttus, Jüri Allik, Anu Realo, Helle Pullmann, Jérôme Rossier, Gregory Zecca, Jennifer Ah-Kion, Denis Amoussou-Yéyé, Martin Bäckström, Rasa Barkauskiene, Oumar Barry, Uma Bhowon, Fredrik Björklund, Aleksandra Bochaver, Konstantin Bochaver, Gideon P. de Bruin, Helena F. Cabrera, Sylvia Xiaohua Chen, A. Timothy Church, and 16 others. 2012. [Comparability of self-reported conscientiousness across 21 countries](#). *European Journal of Personality*, 26(3):303–317.
- John B Nezlek, Astrid Schütz, Michela Schröder-Abé, and C Veronica Smith. 2011. A cross-cultural study of relationships between daily social interaction and the five-factor model of personality. *J. Pers.*, 79(4):811–840.
- OpenAI. 2024. [Gpt-4o mini: advancing cost-efficient intelligence](#).
- Delroy L. Paulhus and Kevin M. Williams. 2002. [The dark triad of personality: Narcissism, machiavellianism, and psychopathy](#). *Journal of Research in Personality*, 36(6):556–563.
- Angela Ramirez, Mamon Alsalihi, Kartik Aggarwal, Cecilia Li, Liren Wu, and Marilyn Walker. 2023. Controlling personality style in dialogue with zero-shot prompt-based learning. *arXiv preprint arXiv:2302.03848*.
- Nils Reimers and Iryna Gurevych. 2019. Sentence-bert: Sentence embeddings using siamese bert-networks. *arXiv preprint arXiv:1908.10084*.
- Yossi Rubner, Carlo Tomasi, and Leonidas J Guibas. 2000. The earth mover’s distance as a metric for image retrieval. *International journal of computer vision*, 40(2):99–121.
- Aadesh Salecha, Molly E. Ireland, Shashanka Subrahmanya, João Sedoc, Lyle H. Ungar, and Johannes C. Eichstaedt. 2024. [Large language models display human-like social desirability biases in big five personality surveys](#). *PNAS Nexus*, 3(12):pgae533.
- 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.
- Gregory Serapio-García, Mustafa Safdari, Clément Crépy, Luning Sun, Stephen Fitz, Peter Romero, Marwa Abdulhai, Aleksandra Faust, and Maja Matarić. 2023. Personality traits in large language models. *arXiv preprint arXiv:2307.00184*.
- Yan Tao, Olga Viberg, Ryan S Baker, and René F Kizilcec. 2024. Cultural bias and cultural alignment of large language models. *PNAS nexus*, 3(9):pgae346.
- An Yang, Baosong Yang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Zhou, Chengpeng Li, Chengyuan Li, Dayiheng Liu, Fei Huang, Guanting Dong, Haoran Wei, Huan Lin, Jialong Tang, Jialin Wang, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Ma, and 43 others. 2024. [Qwen2 technical report](#). *arXiv preprint arXiv:2407.10671*.
- Zheng Zhang, Ning Li, Qi Liu, Rui Li, Weibo Gao, Qingyang Mao, Zhenya Huang, Baosheng Yu, and Dacheng Tao. 2025. The other side of the coin: Exploring fairness in retrieval-augmented generation. *arXiv preprint arXiv:2504.12323*.

A CulturalPersonas: Dataset Generation

To select documents for the RAG model, we select highly cited academic sourced articles from Google scholar using the following keywords: [*country* culture, *country* demographics, *country* norms and practices].⁵

To implement our RAG architecture, we utilize *gpt-4o* and OpenAIEmbedding's *text-embedding-3-large* model to compute embeddings for each country's documents. We use a FAISS vector database to store these embeddings and use similarity search, retrieving the top 5 documents for each query. We experiment with additional values of top *k* documents and identify 5 to represent consistent and diverse dataset creation.

In Figure 6, we present the set of prompts used to generate the original dataset. Additionally, we perform a set of checks after each generation to ensure improved quality of the generation. We thus prompt the model to reflect on its generation and regenerate at each step if necessary. In Figure 7, we present the set of check prompts used.

Additionally, we provide a set of examples from *CulturalPersonas*, 1 from each OCEAN trait:

Trait: Openness

Country: Brazil

Cultural Norm: The concept of "jeitinho" refers to creative ingenuity in rapidly achieving short-term solutions to problems, which is an essential aspect of Brazilian behavior.

Scenario: Lucas gets stuck in a heavy traffic jam and is at risk of missing his important meeting. He decides to weave through the side streets, using shortcuts known only to locals, and arrives a bit early.

Question: Imagine you are Lucas and you've just discovered a colorful street festival happening along one of the shortcuts you're taking. Do you take a few moments to explore the festival and immerse yourself in the local culture, or do you stick strictly to your route to ensure you arrive at the meeting?

Answers:

- a) High: I'll definitely take some time to explore the festival; experiencing local culture is always worth it.
- b) Moderately High: I'll explore the festival quickly, it's a unique chance to enjoy some local culture before moving on.

c) Medium: I might glance around the festival for a brief moment, but I won't delay for too long.

d) Moderately Low: I'll probably continue on my way, perhaps stopping if something really catches my eye.

e) Low: I'll stick to my route; I'm not interested in getting sidetracked by the festival.

Trait: Conscientiousness

Country: India

Cultural Norm: Indian families often live together in the same house, blurring the lines between personal and public spaces.

Scenario: Sana's elder brother and his wife live on the top floor of their family home, but they come downstairs to dine with the rest of the family almost every day.

Question: Assume you are in charge of hosting the family dinners at home. How do you ensure that everything runs smoothly and that everyone is comfortable, including handling the logistics of seating and dietary preferences?

Answers:

- a) High: I create a detailed plan well in advance, considering seating arrangements, dietary restrictions, and decor to ensure a pleasant atmosphere for all.
- b) Moderately High: I make sure to organize the key aspects ahead of time and check in with family members about their seating and dietary needs.
- c) Medium: I plan the dinner with some thought to seating and food preferences, but also adapt as needed on the day.
- d) Moderately Low: I arrange the basic setup and rely on family members to adjust their seating and dietary preferences.
- e) Low: I focus on preparing the meal and let everyone self-organize their seats and dietary choices.

Trait: Extraversion

Country: Saudi Arabia

Cultural Norm: I focus on preparing the meal and let everyone self-organize their seats and dietary choices.

Scenario: Ahmed works as an engineer and spends his days at the office, while his wife, Noor,

⁵We will include the list of documents used for each country upon publication.

takes care of their two children and manages the household chores, including cooking and cleaning.

Question: Imagine Ahmed is invited to a large family gathering, but he has a long working week ahead. Imagine you are Ahmed, how enthusiastic are you about attending the event and engaging with extended family members and not just your wife?

Answers:

- a) High: Very enthusiastic, as attending the gathering is a chance to connect and socialize with everyone, which I look forward to greatly.
- b) Moderately High: Quite enthusiastic, I enjoy meeting family and will find the time to catch up and socialize with everyone I can.
- c) Medium: I'm fairly interested and will attend, though I might not stay the whole time due to my busy schedule.
- d) Moderately Low: Somewhat hesitant; I prefer smaller gatherings and might focus more on resting from work.
- e) Low: I'm not very interested; I'd rather rest and recharge for the upcoming week than socialize at a large gathering.

Trait: Agreeableness

Country: Japan

Cultural Norm: Japanese society highly values punctuality and time management, reflecting their commitment to efficiency and respect for others' time.

Scenario: A Japanese businessman arrives for a meeting ten minutes early, uses the time to review notes, and places his phone on silent to avoid distractions during the meeting.

Question: Assume you are the businessman and during your wait, a junior colleague approaches you for advice. Would you willingly assist them despite the approaching meeting, or prioritize your preparation time?

Answers:

- a) High: I prioritize helping my colleague, understanding that teamwork is essential, confident that I can manage my preparation swiftly afterward without being late.
- b) Moderately High: I assist my colleague with brief, focused guidance, ensuring I still have enough time to complete my preparation.

c) Medium: I offer quick advice but politely inform them I must continue my preparation soon due to the time constraint.

d) Moderately Low: I apologize and explain that I need to prepare, suggesting we meet after the meeting for a more detailed discussion.

e) Low: I inform them that I am unavailable at the moment, emphasizing that my preparation takes priority due to the upcoming meeting.

Trait: Neuroticism

Country: South Africa

Cultural Norm: In South Africa, the concept of Ubuntu, which emphasizes the interconnectedness of humanity, is deeply ingrained in the culture, fostering a sense of community and shared responsibility among individuals.

Scenario: A student at a local school struggles academically, and fellow students form a study group to help them catch up. They share notes, explain difficult concepts, and provide moral support, ensuring their peer feels included and motivated.

Question: Imagine you are feeling overwhelmed with your own academic workload while still helping the struggling student. How do you handle your stress to ensure you can continue to be a supportive member of the study group?

Answers:

- a) High: I find myself worrying frequently and sometimes find it hard to focus, but I try to manage by practicing deep breathing and reminding myself that this phase will pass.
- b) Moderately High: I feel anxious at times, but I take short breaks to clear my head and prioritize tasks to stay on top of everything.
- c) Medium: I acknowledge the stress but focus on planning and managing my schedule effectively to balance both my work and supporting the student.
- d) Moderately Low: I generally stay calm and rely on maintaining a steady study routine to manage my tasks without getting overly stressed.
- e) Low: I rarely feel stressed in these situations, so I remain relaxed and continue to offer my help with a positive outlook.

ok i have issues in this appendix. please help fix:

B CulturalPersonas: Human Validation

In this section, we describe our annotation process to verify the quality of our benchmark. We detail the initial tasks of evaluating: (1) norm relevance, (2) scenario quality, (3) Q&A pair quality. We further detail and provide agreement scores for the norm categorization task.

To validate the quality of our benchmark, we recruit two annotators from each country (undergraduate and masters students) who were born and brought up in the country. Each annotator is asked to label the norms: *Is the norm relevant to the country?*. Then, they are asked to further annotate 10 randomly chosen scenarios: *Which norm does this scenario most relate to?*. Finally, they are asked to evaluate the quality of 50 Q&A pairs: *Does the question truly test the designated OCEAN trait?* and *Do the answers as currently ranked exhibit decreasing levels of the trait?*

To better ground our extracted norms through cultural theory, we ask annotators to complete a secondary norm categorization task. We ask each annotator to map each cultural norm to one of the Hofstede dimensions which include: *Power Distance Index*, *Individualism v. Collectivism*, *Masculinity v. Femininity*, *Uncertainty Avoidance Index*, *Long-Term vs. Short-Term Orientation*, *Indulgence vs. Restraint*, with values ranging from 0 to 100. We ask annotators to map each cultural norm to a Hofstede dimension (or other) and determine whether the cultural norm exhibits high (60-100), med (40-59), or low (0-39) values of the dimension. In Figure 8, we report the agreement scores for each country on (1) norm categorization (2) norm expression. We find that over 80% of the cultural norms are able to fit into Hofstede’s dimensions with (100% of them) showing correct expressions. This shows that our existing RAG architecture is able to generate a diverse set of cultural norms.

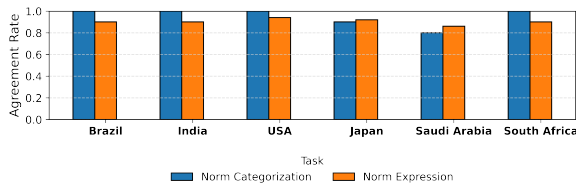


Figure 8: Annotator agreement scores for secondary tasks: (1) Norm categorization into Hofstede dimensions, (2) Norm expression mapped into high, med, low.

Evaluator	O	C	E	A	N	Answer Style
BFI	10	8	8	9	8	Likert-5
IPIP-120	24	24	24	24	24	Likert-5
IPIP-300	60	60	60	60	60	Likert-5
TRAIT	1000	1000	1000	1000	1000	Binary
Big5Chat	106	97	99	98	100	Binary
CulturalPersonas (Ours)	100	100	100	100	100	Likert-5

Table 6: Questions per trait (OCEAN) and answer format across personality evaluators. *CulturalPersonas* consists of 500 total questions per country (total 600 per trait for all 6 countries).

C Personality Evaluators

In this section, we elaborate on the personality evaluation methods used in our experiments. Given that personality is rooted in psychological theory and typically assessed through expert-validated instruments, we begin by selecting three widely recognized tests that have also been applied to LLMs in prior work: the Big Five Inventory (BFI) and two versions of the International Personality Item Pool (IPIP-120 and IPIP-300). Additionally, we incorporate two state-of-the-art personality evaluators specifically developed for language models—TRAIT and Big5Chat. Table 6 provides an overview of the number of questions per trait and the corresponding answer formats. Our benchmark, *CulturalPersonas*, comprises six datasets—one for each selected country—each containing 500 questions tailored to assess culturally grounded personality expression.

D Additional Results

In this section, we present additional analyses and results for personality expression in LLMs. In particular, we focus on (1) Alignment metrics for all 3 models across all 6 benchmarked countries (§ D.1), (2) Relative improvements in personality alignment with trait breakdowns (§ D.1), (3) Open-ended personality expression across various benchmarks and models (§ D.2).

D.1 Personality Alignment Across Various LLMs

We first present all results for multiple choice selection (MCS) and open-ended generation (OEG) split by model in Figures 9 and 10. Our results show that MCS setting is much easier with lower alignment metrics. However, in both settings, *CulturalPersonas* can yield closer alignment with real human populations. GPT-4o-mini is able to achieve the strongest alignment across most countries. Qwen2

however tends to struggle more with expressive and underrepresented countries such as Brazil and South Africa. Llama3 tends to perform on average (in the middle) for most countries, only beating GPT-4o-mini on India.

Next, we present the relative improvement rates for each country using *CulturalPersonas* over both settings: MCS and OEG in Tables 7 and 8. We find that the models generally show strong yields for *Conscientiousness* and *Openness*, with more moderate to low gains for traits such as *Neuroticism* and *Agreeableness*. In particular, gains are most pronounced in East Asian countries i.e. Japan, India, and Saudi Arabia and Africa i.e. South Africa.

Country	TRAIT					Big5Chat				
	O	C	E	A	N	O	C	E	A	N
BR	6.4	7.2	3.4	4.8	5.2	16.4	17.5	10.8	11.3	9.0
IN	13.4	15.6	9.1	8.6	11.3	11.4	10.4	6.7	6.4	8.6
JP	20.4	17.5	14.5	14.2	20.4	6.8	7.2	6.9	4.7	5.4
US	8.8	7.9	3.4	2.3	5.0	7.0	6.6	3.0	2.5	4.2
SA	13.8	12.5	3.0	2.0	7.0	15.2	14.3	7.5	5.5	10.0
ZA	12.0	11.6	3.5	3.4	8.0	15.0	13.9	4.0	2.0	8.5

Table 8: Relative improvement (%) in Wasserstein distance using *CulturalPersonas* across OEG settings for each OCEAN trait

D.2 Analysis of Open-Ended Generation Results

In this section, we provide additional examples on various models express personality in the different benchmarked countries. In Table 9, we present examples on how GPT-4o-mini provides more culturally appropriate responses when prompted with *CulturalPersonas* compared to other existing tools.

Additionally, we show examples of personality expression across various models. In Tables 10 and 11, we show variations of how *Extraversion* and *Agreeableness* appear in Qwen2 v. Llama3 v. GPT-4o-mini. Qwen2 often struggles with expressing *Agreeableness* in countries e.g. Japan. The example highlights more Western-centric norms in its answers compared to GPT-4o-mini and Llama3 whose responses highlight more of a passive and indirect method for maintaining harmony. On the other hand, Llama3 struggles with expressing *Extraversion* in expressive nations such as Brazil. We can see when comparing with model responses from Qwen2 and GPT-4o-mini, these responses are much more expressive and engaging in ways related to Brazilian norms.

E Additional Ablations

In this section, we present ablation experiments by removing our selected persona prompts i.e. removing cultural priming and observing distributional alignment changes between models and human data § E.1. We also present results from using a simple prompt with only a country persona in § E.2.

E.1 Cultural Priming Ablations

In this section, we provide results from all our cultural priming experiments (see Figure 11). We present deviations in performance (Wasserstein distances) for each country for each benchmarked model. We find that cultural priming indeed yields positive performance gains across all countries. Although priming for the USA provides the lowest gains, incorporating cultural norms still yields small, measurable improvements.

E.2 Simple Prompting

In this section, we provide results from our simple prompting experiments. Building on prior findings that simple persona prompts often lack sufficient context to elicit demographic-specific responses (Tao et al., 2024; Ghosh et al., 2024), we examine the extent to which cultural priming offers more effective guidance. In Table 12, we present the average Wasserstein distances for each country given a simple persona prompt i.e. “You are someone from {country}.”

Country	BFI					IPIP-120					IPIP-300					TRAIT					Big5Chat				
	O	C	E	A	N	O	C	E	A	N	O	C	E	A	N	O	C	E	A	N	O	C	E	A	N
BR	19.9	19.8	31.3	2.4	22.2	25.1	24.9	2.3	22.8	47.5	21.3	20.8	11.1	18.9	28.9	27.4	26.9	13.5	34.4	28.4	23.1	23.0	23.5	19.1	21.4
IN	16.0	15.7	7.5	34.6	0.2	17.4	16.5	22.0	6.4	17.7	16.0	15.2	3.0	32.2	6.7	27.6	27.4	20.0	23.2	35.8	24.8	24.4	31.0	19.6	18.7
JP	17.4	16.4	29.2	1.6	15.6	24.3	23.6	23.5	42.7	0.8	28.2	28.5	52.3	5.9	23.6	35.1	34.4	87.7	9.3	3.6	33.6	33.2	12.1	73.8	11.3
US	13.4	12.9	10.2	5.8	18.7	18.6	17.6	38.5	9.8	1.6	19.1	19.2	41.6	1.9	9.7	17.3	17.6	23.6	1.5	23.9	27.1	26.6	4.5	18.3	54.7
SA	22.5	21.9	30.7	14.3	16.2	25.7	25.4	5.8	31.4	36.8	24.5	24.3	31.8	27.0	10.5	25.6	25.2	47.5	22.4	0.8	25.1	25.2	11.4	1.9	58.4
ZA	26.7	26.3	28.0	8.4	40.0	23.1	23.3	0.6	58.8	7.2	20.9	20.3	5.2	52.3	0.4	31.4	30.4	26.5	3.1	58.7	34.1	33.3	17.6	56.7	22.9

Table 7: Relative improvement (%) in Wasserstein distance using *CulturalPersonas* across MCS settings for each OCEAN trait

Country	Benchmark	Response
Japan	TRAIT	I might profusely apologize to my colleague.
	Big5Chat	I am very sorry. I did not intend to hurt you.
	CulturalPersonas	I would bow deeply and express regret for hurting my colleague.
USA	TRAIT	I would say sorry and explain it was unintentional.
	Big5Chat	I apologize and take responsibility to make it right.
	CulturalPersonas	I would acknowledge my mistake and offer to fix it.
Brazil	TRAIT	I'd approach them and try to talk it out.
	Big5Chat	Sorry! Let's chat to smooth things over.
	CulturalPersonas	I would show concern and invite a conversation.
India	TRAIT	I would offer to help with their tasks.
	Big5Chat	I'm sorry. I didn't mean to upset you. Can I make it up?
	CulturalPersonas	I would acknowledge their hurt and apologize respectfully.
Saudi Arabia	TRAIT	I'd apologize and hope to maintain harmony.
	Big5Chat	My apologies, I didn't intend offense.
	CulturalPersonas	I would sincerely apologize and show respect.
South Africa	TRAIT	I'd apologize and ask how to make things right.
	Big5Chat	I'm sorry. Let me know how I can help.
	CulturalPersonas	I would take responsibility and seek a positive way forward.

Table 9: Illustrative examples of culturally appropriate personality expressions (GPT-4o-mini) across six countries for a workplace scenario testing *Agreeableness*. *CulturalPersonas* often provides culturally rich responses such as "bowing" in Japan, "respecting" in Saudi Arabia and India, and "showing concern" in Brazil.

Scenario: You are at a friend's birthday party, and you only know a few people there. What do you do?		
Llama3	GPT-4o-mini	Qwen2
I'd not stick close to my friends, but I might feel shy even though I'd try to join into the groups.	I'd smile, start conversations with people nearby, and try to join in a group activity or dance.	I'd compliment someone's outfit or joke about the music to break the ice, then gradually join the group.

Table 10: Cultural expressions of *Extraversion* in a Brazil-based social scenario. While all answers show moderate *Extraversion*, Llama3 demonstrates a more indirect approach. GPT-4o-mini and Qwen2 align more with Brazilian sociability norms.

Scenario: Your manager assigns extra work to you at the last minute, even though it wasn't your responsibility. How do you respond?		
Llama3	GPT-4o-mini	Qwen2
I would accept the task without complaint and do my best to complete it, showing commitment to team harmony.	I'd politely accept the work and later mention, with deference, that I'm happy to help this time but may struggle if this becomes frequent.	I'd agree to take it on and say, "No worries, I've got it," while gently suggesting we plan better next time to avoid last-minute stress.

Table 11: Cultural expressions of *Agreeableness* in a Japan-based workplace scenario. Llama3 and GPT-4o-mini reflect Japanese-style cooperation, while Qwen2 uses direct yet friendly feedback, aligned with US norms.

Country	BFI	IPIP-120	IPIP-300	TRAIT	Big5Chat	CulturalPersonas
Brazil	19.12	24.52	20.20	26.12	22.02	21.44
India	14.80	16.00	14.62	26.80	23.70	19.23
Japan	16.04	22.98	27.70	34.02	32.80	28.34
USA	12.20	17.22	18.30	16.78	26.24	31.34
Saudi Arabia	21.12	25.02	23.62	24.30	24.40	12.45
South Africa	25.88	22.60	19.82	30.02	32.92	20.23

Table 12: Average relative improvement (%) in Wasserstein distance averaged across OCEAN traits using cultural priming over standard country context prompt for each country and evaluation set.

Prompts used for CulturePersonas Dataset Generation

Cultural Norm Prompt:

"Provide me 10 unique sentences highlighting the core values/important aspects of individuals living in {country}."

Scenario Prompt:

Please generate a detailed scenario for the following country: {country} and cultural norm: {norm}. Please keep the scenario within 1-2 sentences.

Question Prompt:

You are an expert at analyzing user's personality types, specifically their Big5 attributes which include (Openness to Experience, Conscientiousness, Extraversion, Agreeableness, and Neuroticism). You will be given a scenario and you will be required to generate 5 multiple choice questions to test each of the Big5 traits. Each question should test exactly one of the traits and each trait should be tested exactly once.

Please generate the questions so that they are role play scenarios where the user can imagine themselves in the given scenario. You should be using keywords such as "imagine you are ..." or "assume you are in the position of..."

Please generate creative questions that can have open-ended action-style responses where different users may respond to the situation in different ways depending on their personality. Now, please generate questions based on the following scenario: {scenario} for country: {country}. Please keep each question within 2-3 sentences.

Answer Prompt:

We are creating a dataset to test a user's personality traits from country: {country}.

You will be provided a question along with the Big5 trait (Openness to Experience, Conscientiousness, Extraversion, Agreeableness, and Neuroticism) it is meant to test.

You are required to generate 5 different multiple choice answers.

The first answer should express very high levels of the trait. The second answer should express moderately high levels of the trait. The third answer should express medium levels of the trait. The fourth answer should express moderately low levels of the trait. The fifth answer should express low levels of the trait.

Please make sure no answers are disrespectful and are viable responses to the situation in the question. Do not exhibit obviously low or high levels in the answers.

Please limit the answers of the question to 1-2 sentences.

Please generate 5 multiple choice answers to test the trait: {trait} for the following question: {question}.

Figure 6: Multi-step prompts used to generate CulturePersonas Dataset. We generate 500 total questions for each country: 100 question and answer sets for each OCEAN trait.

Reflection Prompts used for CulturePersonas Dataset Re-Generation

Scenario Check Prompt:

Country: {country} has the following cultural norm: {norm}.
Please check if the following scenario: {scenario} corresponds with the country's norm.
Please first provide a "Yes" or "No" and then provide any reasons.

Scenario Regeneration Prompt:

Please regenerate a scenario for the following country: {country} and cultural norm: {norm}.
This is the previous scenario you generated: {scenario} and here is the feedback you received: {feedback}.
Please only generate the scenario and do not provide any explanations.

Question Check Prompt:

Here is a possible scenario in country {country}: {scenario}.
To test a user's Big5 {trait} levels, here is a proposed question to ask them: {question}.
Please check if this question corresponds with the scenario and tests the trait: {trait}.
Also, please make sure the question is a role-play style question where the users can imagine themselves in the situation. For example, the question should include keywords like "imagine you are..."
Finally, make sure the question is creative and can have various open-ended action-style responses.
Please first provide a "Yes" or "No" and then provide any reasons.

Question Regeneration Prompt:

You are an expert at analyzing user's personality types from {country}, specifically their Big5 attributes which include (Openness to Experience, Conscientiousness, Extraversion, Agreeableness, and Neuroticism).
Please generate a question to test a user's {trait} levels given this scenario: {scenario}.
Please generate the question so that they are role play scenarios where the user can imagine themselves in the given scenario. You should be using keywords such as "imagine you are ..." or "assume you are in the position of...".
Please generate a creative question (2-3 sentences) that can have open-ended action-style responses where different users may respond to the situation in different ways depending on their personality.
This is the previous question you generated: {question} and the feedback you received: {feedback}.

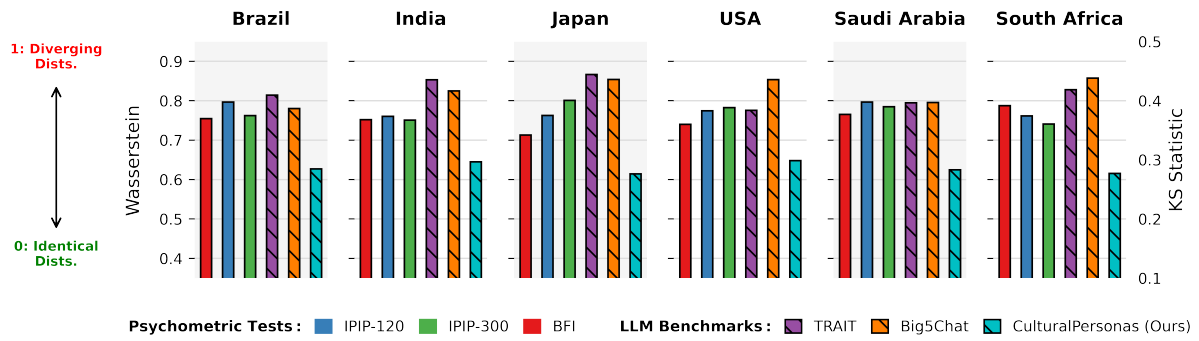
Answer Check Prompt:

We are creating a dataset to test a user's personality traits from country: {country}.
You will be provided a question along with the Big5 trait (Openness to Experience, Conscientiousness, Extraversion, Agreeableness, and Neuroticism) it is meant to test.
You are required to generate 5 different multiple choice answers.
The first answer should express very high levels of the trait. The second answer should express moderately high levels of the trait. The third answer should express medium levels of the trait. The fourth answer should express moderately low levels of the trait. The fifth answer should express low levels of the trait.
Please make sure no answers are disrespectful and are viable responses to the situation in the question. Do not exhibit obviously low or high levels in the answers.
Please limit the answers of the question to 1-2 sentences.
Please generate 5 multiple choice answers to test the trait: {trait} for the following question: {question}.

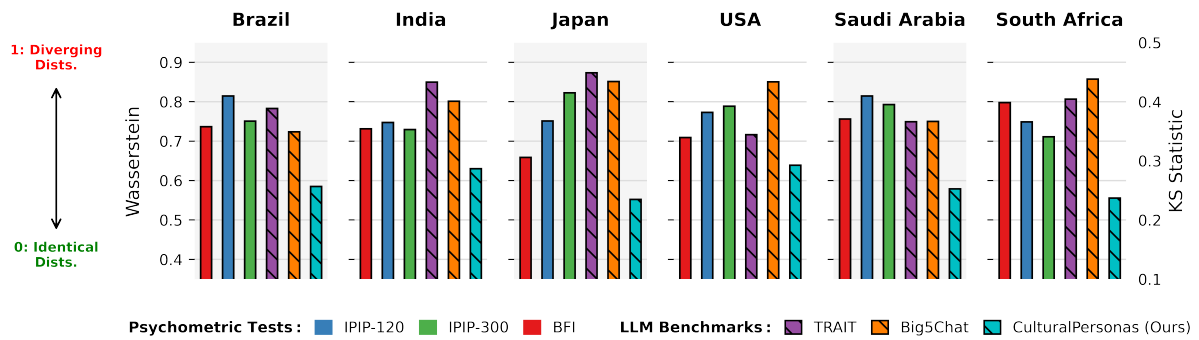
Answer Regeneration Prompt:

We are creating a dataset to test a user's personality traits from country: {country}.
You will be provided a question along with the Big5 trait (Openness to Experience, Conscientiousness, Extraversion, Agreeableness, and Neuroticism) it is meant to test.
You are required to generate 5 different multiple choice answers that exhibit {trait}.
The first answer should express very high levels of the trait. The second answer should express moderate levels of the trait. The third answer should express medium to low levels of the trait. The fourth answer should express moderately low levels of the trait. The fifth answer should express low levels of the trait.
Please make sure no answers are disrespectful and are viable responses to the situation in the question. Do not exhibit obviously low or high levels in the answers.
Please limit the answers of the question to 1-2 sentences.
This is the question: {question}. You previously generated these answers --> answer_high: {answer_high}, answer_moderately_high: {answer_moderately_high}, answer_medium: {answer_medium}, answer_moderately_low: {answer_moderately_low} and answer_low: {answer_low}. And you received this feedback: {feedback}. Please improve the answers.

Figure 7: At each step of the dataset generation pipeline, the model self-reflects on the content generation and regenerates content if necessary.

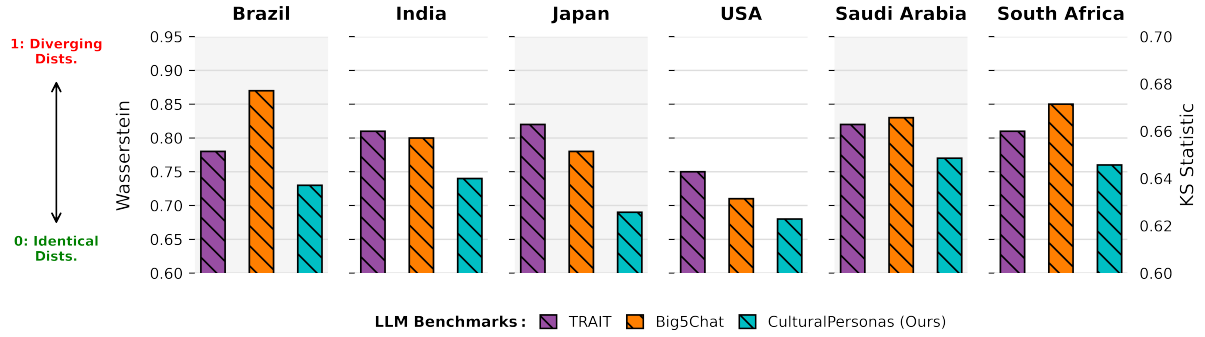


(a) MCS results from Llama3

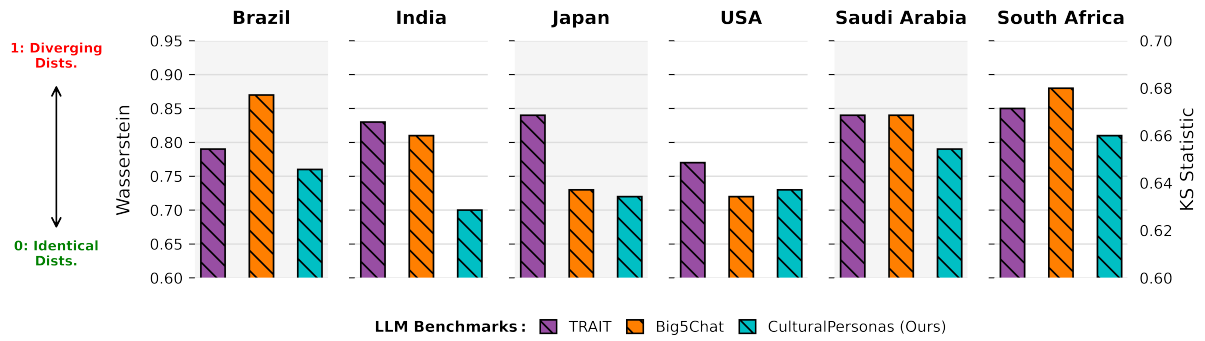


(b) MCS results from Qwen2

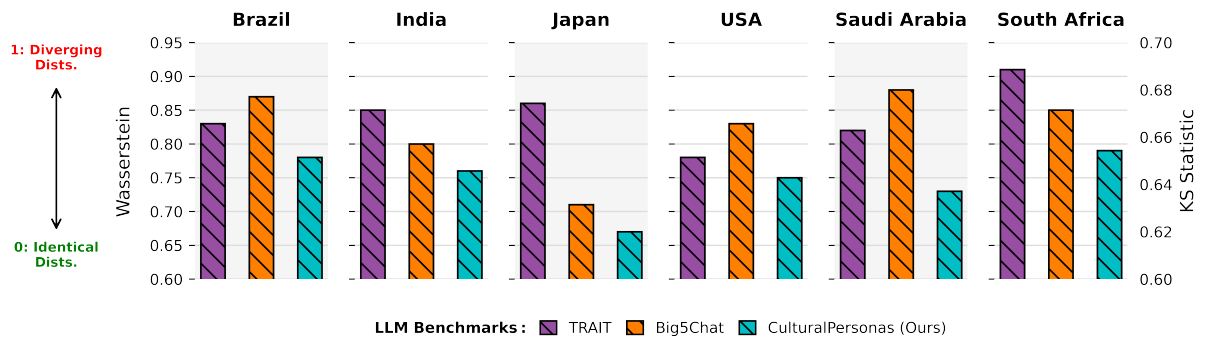
Figure 9: Cross-cultural Wasserstein alignment in Multiple-Choice Selection (MCS) for Llama3 and Qwen2. Lower scores indicate closer alignment with human personality distributions. Overall, we observe *CulturalPersonas* provides higher alignment across different countries.



(a) GPT-4o-mini

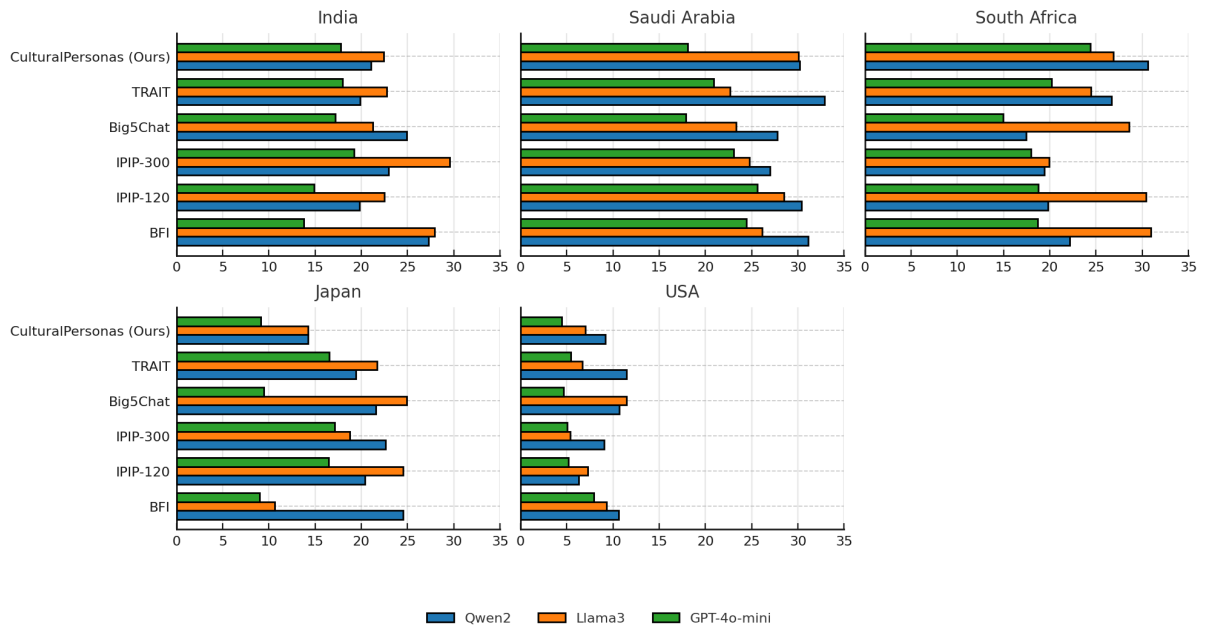


(b) Llama3

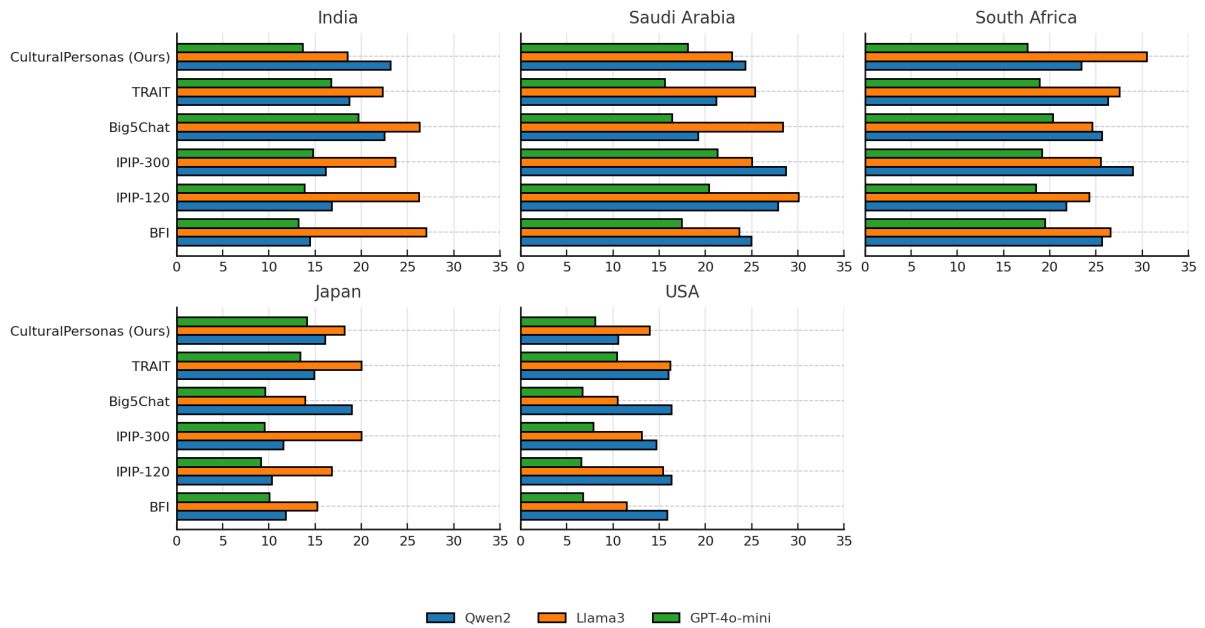


(c) Qwen2

Figure 10: Cross-cultural Wasserstein alignment in Open-Ended Generation (OEG) for all models. Lower scores indicate closer alignment with human personality distributions. We report higher scores in OEG than MCS, and models tend to diverge in performance much more. Although GPT-4o-mini yields best results overall, Llama3 tends to yield better results for India. Qwen2 tends to struggle with more diverse countries such as Brazil and South Africa.



(a) Cultural Priming Ablations: MCS



(b) Cultural Priming Ablations: OEG

Figure 11: Results from cultural priming ablations for each country and model for both MCS and OEG settings. We observe large increases (%) in alignment (Wasserstein distances) for various countries including underrepresented and highly expressive nations.