

CAMELLIA : BENCHMARKING CULTURAL BIASES IN LLMs FOR ASIAN LANGUAGES

Tarek Naous¹, Anagha Savit¹, Carlos Rafael Catalan², Geyang Guo¹, Jaehyeok Lee³, Kyungdon Lee³, Lheane Marie Dizon², Mengyu Ye⁴, Neel Kothari¹, Sahajpreet Singh⁵, Sarah Masud⁶, Tanish Patwa¹, Trung Tanh Tran⁷, Zohaib Khan⁸, Alan Ritter¹, JinYeong Bak³, Keisuke Sakaguchi⁴, Tanmoy Chakraborty⁹, Yuki Arase¹⁰, Wei Xu¹

¹Georgia Institute of Technology, ²Samsung R&D Institute Philippines, ³Sungkyunkwan University, ⁴Tohoku University, ⁵National University of Singapore, ⁶University of Copenhagen, ⁷Takenote.ai, ⁸University of Michigan, ⁹Indian Institute of Technology Delhi, ¹⁰Institute of Science Tokyo

✉ tareknaous@gatech.edu

🌐 tareknaous/camellia

ABSTRACT

As Large Language Models (LLMs) gain stronger multilingual capabilities, their ability to handle culturally diverse entities becomes crucial. Prior work has shown that LLMs often favor Western-associated entities in Arabic, raising concerns about cultural fairness. Due to the lack of multilingual benchmarks, it remains unclear if such biases also manifest in different non-Western languages. In this paper, we introduce Camellia, a benchmark for measuring entity-centric cultural biases in nine Asian languages spanning six distinct Asian cultures. Camellia includes 19,530 entities manually annotated for association with the specific Asian or Western culture, as well as 2,173 naturally occurring masked contexts for entities derived from social media posts. Using Camellia, we evaluate cultural biases in four recent multilingual LLM families across various tasks such as cultural context adaptation, sentiment association, and entity extractive QA. Our analyses show a struggle by LLMs at cultural adaptation in all Asian languages, with performance differing across models developed in regions with varying access to culturally-relevant data. We further observe that different LLM families hold their distinct biases, differing in how they associate cultures with particular sentiments. Lastly, we find that LLMs struggle with context understanding in Asian languages, creating performance gaps between cultures in entity extraction.

1 INTRODUCTION

Large Language Models (LLMs) have rapidly integrated into modern technology, serving users from diverse cultures (Adilazuarda et al., 2024). Among the vast range of text they process, LLMs frequently encounter entities such as people’s names, locations, or food dishes, which are pervasive in text corpora (Wolfe & Caliskan, 2021; Pawar et al., 2025a) and often appear in user prompts (Li et al., 2024a; Wang et al., 2025). Importantly, entities carry cultural associations, making it essential for LLMs to handle culturally diverse entities fairly. However, past work has shown that these cultural associations can significantly influence LLMs, leading to biased behaviors (An et al., 2024; Wan et al., 2023). The recent study of Naous et al. (2024) demonstrated how such biases manifest when testing LLMs in Arabic, where models showed better performance on entities associated with Western culture compared to those linked to Arab culture. A natural question is *whether similar LLM cultural biases would also manifest in other non-Western languages*.

To this end, we introduce Camellia (Cultural Appropriateness Measure Set for LLMs in Asian Languages), a benchmark for measuring entity-centric cultural biases in 9 non-Western languages spoken in the Asian continent: Chinese (zh), Japanese (ja), Korean (ko), Vietnamese (vi), Urdu

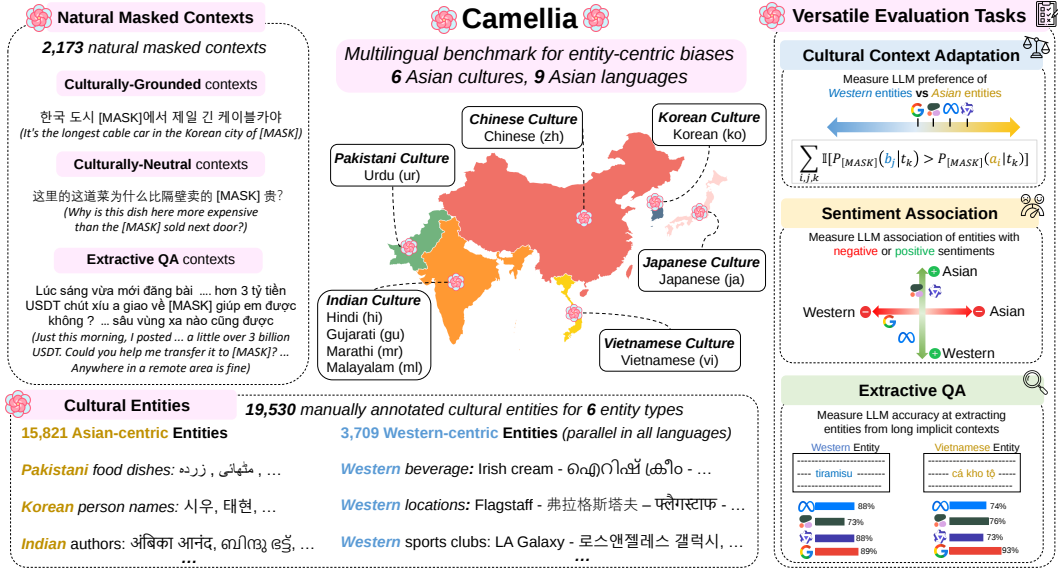


Figure 1: We construct Camellia, a benchmark to measure cultural biases for six Asian cultures, covering nine languages. Camellia provides 2,173 naturally-occurring masked contexts categorized into: culturally-grounded, culturally-neutral, and extractive QA. Camellia also provides 19,530 culturally relevant entities that contrast the respective Asian cultures vs. Western culture across six different entity types that exhibit cultural variation. The masked contexts and entities in Camellia enable the measurement of cultural biases in LLMs via versatile task setups.

(ur), Hindi (hi), Malayalam (ml), Marathi (mr), and Gujarati (gu), covering 6 distinct cultures in Asia (see Figure 1). Following the data curation process outlined in CAMEL (Naous et al., 2024), we undertook a year-long collaboration with native speakers to collect and annotate 19,530 cultural entities across six entity types contrasting Asian and Western cultures (§2.1). We also curate 2,173 naturally occurring masked contexts for entities spanning all nine languages (§2.2). Moreover, we provide English translations for each entity and masked context in Camellia, enabling direct cross-lingual comparisons for testing LLMs in English vs the respective Asian language.

Using Camellia, we examine cultural biases in four recent multilingual LLM families (Llama, Qwen, Aya, Gemma) across diverse evaluation setups (§3). Our experiments show how **LLMs can struggle to adapt to Asian cultural contexts in all languages**, assigning higher likelihood for Western-associated entities in 30-40% of cases, even when inappropriate to the context (§3.1). Further, we find that **different model families display their own distinct biases**. When analyzing cultural sentiment associations in LLMs, Qwen shows a higher tendency of associating Asian entities with positive sentiment compared to Western entities, whereas the Llama and Gemma models show the opposite trend (§3.2). Lastly, we show how **LLMs still lack the ability to efficiently grasp context in Asian languages, impacting their cultural fairness in entity extraction**. When tasked with extracting entities from paragraphs, we observed large accuracy gaps in LLMs when entities in the same text were associated with different cultures. In contrast, these gaps were minimal when testing LLMs on the English translations of contexts and entities, where performance is stable regardless of an entity’s cultural association (§3.3).

2 CONSTRUCTING CAMELLIA

This section describes the process of constructing the Camellia benchmark. First, we outline our methodology for collecting culturally-relevant entities across nine different Asian languages (§2.1). We then describe how we collect naturally-occurring masked contexts for entities, which enable testing for entity-centric cultural biases in LLMs across versatile setups (§2.2).

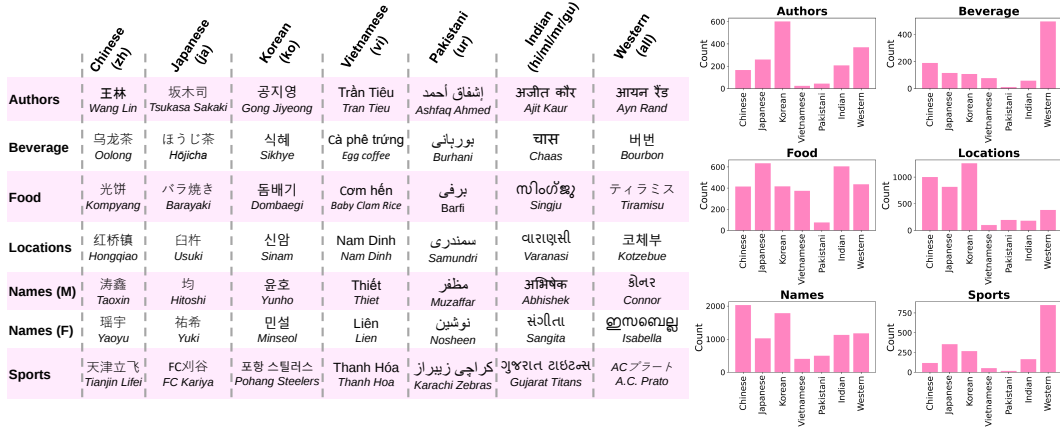


Figure 2: Example per entity type and statistics of respective Asian entities per culture and Western entities in Camellia. Western entities are parallel for all 9 languages while Indian entities are parallel in all Indian languages (§2.1). Camellia also provides an English translation for each entity.

2.1 COLLECTING CULTURAL ENTITIES

Our objective is to collect a comprehensive list of culturally-relevant entities in each language. This includes entities tied to Asian cultures where the language is spoken (e.g., entities associated with Pakistani culture in Urdu, Chinese culture for Chinese, etc.) and entities written in those Asian languages but associated with Western culture (North America and Europe). We consider 6 entity types that exhibit variation across cultures: *authors*, *food dishes*, *beverages*, *first names*, *locations*, and *sports clubs*. To collect entities, we follow the procedure described in the CAMEL benchmark (Naous et al., 2024), which leverages the multilingual Wikidata knowledge base and performs pattern-based extraction on web-crawled data. Figure 2 shows the statistics of Asian-centric and Western entities that we collect and annotate for each language in Camellia.

Extracting Entities from Wikidata. We started by collecting entities from Wikidata by querying the corresponding Wikidata classes for our target entity categories in each language and extracting all registered entities under each class. We found the coverage in Wikidata to be generally sufficient for *authors*, *locations*, and *sports clubs* for all languages. However, the coverage for the other entity types (*food dishes*, *beverages*, *names*) was much less extensive and varied by language. As of 2024, we observed that higher-resource languages had a sizable amount of entities in Wikidata (e.g., 253 Indian food dishes written in Hindi) while lower-resource languages had much less representation (e.g., only 24 Indian food dishes in Malayalam, 37 Pakistani names in Urdu, etc.).

Pattern-based Extraction from Web-Crawls. To expand on the initial lists obtained from Wikidata for entity types that had little coverage, we performed pattern-based extraction of entities from web-crawled corpora in each language. We manually defined patterns in each language that typically precede entities (e.g., *brother/sister named ____* for first names, *recipe of ____* for food dishes, etc.). Using the patterns, we scanned through each language’s partition in the mC4 web-crawl corpus (Xue et al., 2021) and extracted unigrams and bigrams that appeared after a detected pattern. We also accounted for gender inflections if required. This resulted in 5k-10k extractions in each type and language, which were then manually filtered to remove irrelevant extractions and select culturally-relevant entities. Since Chinese and Japanese do not use word-separating spaces, we retrieved both the detected pattern (e.g., “喝”, which means “to drink”) and up to ten surrounding characters in these languages, and then prompted GPT-4o-mini to extract the entity from the captured characters, if any were mentioned. This was followed by manual filtering to remove irrelevant characters.

Annotation by Native Speakers. For each of the 9 Asian languages, one of our authors who is a native speaker manually filtered the extractions from Wikidata and mC4 to identify culturally-relevant entities and remove irrelevant extractions. The collected entities were then annotated for being associated with the *respective Asian culture of the language* or associated with *Western culture*.

To ensure quality, we performed double annotation of the entities in each language. The second annotators consisted of undergraduate or master’s students hired for zh, ja, ko, hi, ml, mr, and gu; and native speaker volunteers for vi and ur. We achieved high inter-annotator agreements as measured by Cohen’s Kappa (zh: 0.85, ja: 0.78, ko: 0.92, vi: 0.80, ur: 0.88, hi: 0.94, ml: 0.83, mr: 0.93, gu: 0.97). The disagreements were then resolved in an adjudication step to decide the final label. We report the detailed annotation guidelines for each entity type in Appendix A.

Translating Entities to English. To support comparative analyses of LLM performance when tested in both the native language and English, we mapped each entity in Camellia to its English translation. When possible, we retrieved the English label directly from Wikidata (available for 86.58% of Wikidata-sourced entities). For entities without an English label and ones extracted from mC4, we manually searched for their most commonly used English transliterated form found online, ensuring that the translations reflect how entities appear in real-world usage.

Parallelizing Western Entities. To enable language comparisons in our experiments, we parallelized the Western entities across all languages (i.e., each entity has a written version in every language). For *authors*, *locations*, and *sports clubs*, we constructed their parallel Western sets directly from Wikidata by extracting the entities of each Western country (North America and Europe) that had a written form in at least 6 of the languages. A lot of these Western entities did not have written versions in Wikidata in low-resource languages (ur, ml, gu, and mr). For those languages, we manually filled in their missing translations.

For the other types of *food*, *beverage*, and *names*, Western entities were collected independently in each language via pattern-based extractions mC4. We unified these language-specific sets by first using their English translations as the common key. Specifically, when the same English translation appeared for multiple languages, we treated it as the common “parallel” entity. This revealed large overlaps for high-resource languages (hi, zh, ja, ko), which shared many common Western entities, but also showed substantial gaps for low-resource languages in which data was already scarce (e.g., 1k–1.5k food entities needed to be translated to ur). To minimize translation effort while ensuring quality, we randomly sampled 500 unified entities per type and, with the help of annotators, manually completed the missing entries by translating them from English into their languages.

Parallelizing Entities in Indian Languages. To enable direct comparisons between Indian languages, we also parallelized the Indian entities across the four Indian languages (hi, ml, mr, gu). Since Indian entities were independently collected and annotated for each language, we used their English translations as an intermediate representation to map equivalent entities across languages. Annotators then manually translated the missing gaps from English. The majority of Indian cultural entities were initially collected in hi, being the most resource-rich Indian language. In contrast, manual translation efforts were mostly required to map entities into ml, mr, and gu.

2.2 COLLECTING NATURAL MASKED CONTEXTS

To evaluate whether LLMs can distinguish between entities associated with each Asian culture vs. those associated with Western cultures, Camellia provides 2,173 naturally-occurring masked contexts for entities derived from natural discussions by native speakers on X (formerly Twitter).

Following CAMEL (Naous et al., 2024), we collected short contexts that are uniquely suited for the entities associated with each Asian culture, enabling us to assess LLM cultural adaptation. We also collected neutral contexts where entities from any culture were appropriate, helping determine the default inclinations of models in the absence of clear cultural cues. Additionally, we constructed longer contexts that reference entities more implicitly, presenting a challenging setup for testing models at entity identification in an extractive QA format. Accordingly, the masked contexts in Camellia are split into three types: (1) culturally-grounded (Camellia-Grounded), (2) culturally-neutral (Camellia-Neutral), and (3) extractive QA contexts (Camellia-QA).

Contexts for Evaluating Cultural Adaptation. To construct Camellia-Grounded, we searched X using two types of search queries: randomly sampled Asian entities (e.g., [Indian entity], [Japanese entity]), and manually designed patterns that mention a culturally-relevant entity (e.g., the [Chinese] city of, the [Indian] dish, etc.). We then manually inspected the retrieved tweets to

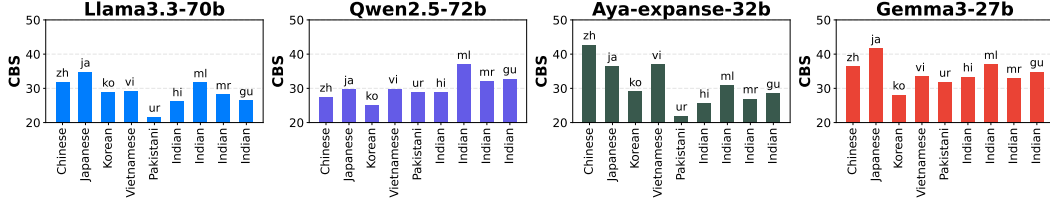


Figure 3: Average Cultural Bias Score (CBS) ($\downarrow$) across entity types achieved by LLMs on culturally-grounded contexts (Camellia-Grounded) for each Asian language. LLMs can struggle to generate the appropriate Asian entities in each culture, assigning better likelihood to Western entities 30-40% of the time. See results per entity type in Appendix C.1.

identify ones that provide suitable cultural contexts. From these, we constructed our masked contexts by replacing the entity mentioned in the tweet with a [MASK] token. Similarly, to construct neutral contexts (Camellia-Neutral), we identified tweets where entities from any culture would be appropriate as [MASK]. Further, we annotated each context with one of three sentiment labels: *positive*, *negative*, or *neutral*. This helps evaluate whether substituting the [MASK] token with the respective Asian or Western entities changes the sentiment predicted by LLMs (§3.2).

Contexts for Extractive QA. In addition to the contexts used to evaluate cultural adaptation in LLMs, we constructed longer, paragraph-level contexts in which entities are mentioned implicitly. These longer contexts enable a challenging evaluation setup for entity extraction, as they require understanding the underlying context to identify the entity. We follow the same keyword search strategy to identify such contexts on the X platform, and replace the mentioned entity with the [MASK] token. Camellia-QA provides ~ 8 -10 of such contexts for each entity type in each language.

Parallelizing Indian Contexts. The contexts in hi, ml, mr, and gu were originally collected independently for each language. To enable comparisons across these Indian languages, we parallelized them by first translating the contexts into English and then into the other Indian languages.

3 ARE CULTURAL BIASES CONSISTENT ACROSS LANGUAGES AND LLMs?

We leverage the cultural entities and masked contexts in Camellia to investigate whether cultural biases in LLMs are persistent across languages and LLMs. We experiment with four recent LLMs with multilingual capabilities: **Llama3.3-70b** (Grattafiori et al., 2024), **Qwen2.5-72b** (Yang et al., 2025), **Aya-expense-32b** (Dang et al., 2024), and **Gemma3-27b** (Team et al., 2025). We test LLMs in three setups: cultural adaptation (§3.1), sentiment association (§3.2), and extractive QA (§3.3).

3.1 CULTURAL CONTEXT ADAPTATION

We first analyze the ability of LLMs to adapt to different Asian cultural contexts by analyzing their assigned likelihood for the respective Asian vs Western entities as [MASK] token fillings.

Cultural Bias Score (CBS). We use the CBS designed by Naous et al. (2024) to measure the level of Western bias in an LLM_θ . CBS is a likelihood-based measure that computes the percentage of an LLM’s preference for Western entities over Asian ones within the same cultural context. Given an entity type D , two type-specific sets of respective Asian entities $A = \{a_i\}_{i=1}^N$ and Western entities $B = \{b_j\}_{j=1}^M$, and a masked context c_k , we compute $CBS_D(LLM_\theta, A, B, c_k)$ per language as:

$$\frac{1}{N \times M} \sum_{i=1}^N \sum_{j=1}^M \mathbb{1}[P_{[MASK]}(b_j|c_k) > P_{[MASK]}(a_i|c_k)], \quad (1)$$

where $P_{[MASK]}$ is the LLM’s probability of an entity filling the [MASK] token. For entities tokenized into multiple tokens, we take the product of the conditional probabilities of each token. For a set

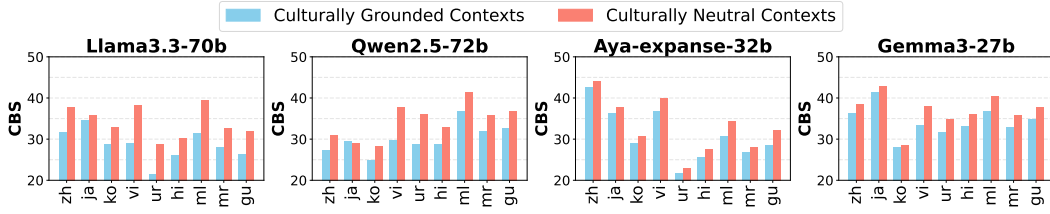


Figure 4: Average CBS across entity types on culturally-grounded contexts vs culturally-neutral contexts. LLMs show more preference towards Western entities in culturally-neutral contexts (higher CBS). CBS scores are lower in culturally-grounded contexts, yet remain close to the neutral case.

of prompts $C = \{c_k\}_{k=1}^K$, the CBS per entity type for an LLM is computed by averaging over all $c_k \in C$. An LLM is considered more Western-biased as its CBS gets close to 100%.

Results. Figure 3 shows the average CBS across entity types achieved on the culturally-grounded contexts of each culture when tested in each language. We observe the following key insights:

LLMs can struggle to distinguish Asian vs. Western entities in Asian languages. Since the contexts we test on are grounded in each Asian culture (only entities associated with the specific Asian culture are appropriate for filling the [MASK]), the CBS is expected to be low (closer to the 0-5% range). However, in most cases, we observe the CBS to be in the 30-40% range. This highlights many situations where LLMs struggle to differentiate between Asian and Western entities, assigning a better likelihood to Western entities despite being inappropriate to the context.

Are models sensitive to cultural grounding? We further analyze if performance changes when testing on the contexts that are culturally neutral (i.e., any entity is an appropriate [MASK] filling in the context). The results are summarized in Figure 4, which shows that CBS scores are higher when contexts are neutral, with LLMs becoming more likely to generate Western entities. However, in the majority of cases, the scores still remain very close to when contexts are culturally grounded. This suggests a lack of sensitivity to cultural contexts in LLMs, whereby their ability to select the appropriate entities at generation time is not greatly impacted by cultural grounding.

Adaptation performance can vary by LLM family. Noticeable differences can be seen in the performance of LLM families developed in different regions. Specifically, we find that the Qwen2.5-72b model that is developed by China-based Alibaba performs the best on Chinese, Japanese, and Korean, compared to the rest of the models. Such a gap likely reflects more access to culturally relevant pre-training data in those languages, enabling the model to learn cultural associations that others would miss. This highlights the importance of data provenance in shaping the cultural competence of LLMs. Moreover, this corroborates the results of past work that shows a better ability of Qwen models at answering questions specific to Chinese culture (Guo et al., 2025).

Adaptation ability for the same culture can vary by resource availability. In the Indian setting, performance varied based on the resource availability of languages. Models performed relatively better when tested in Hindi but struggled more when tested in lower-resource languages as Malayalam, Marathi, and Gujarati. Notably, this trend is consistent across all models, reflecting similar access to training data proportions for those languages. In practice, this makes the adaptation ability of LLMs to Indian contexts skewed towards Hindi, privileging one linguistic community over others.

3.2 SENTIMENT ASSOCIATION

We examine whether LLMs subtly associate entities from Asian or Western cultures with specific sentiments by analyzing their behavior on sentiment analysis.

Setup. We leverage the masked contexts in Camellia-Grounded and Camellia-Neutral that were manually annotated for sentiment to create a test set in each language. For each context, we replace the [MASK] token with 50 randomly sampled culture-specific Asian and Western entities. This

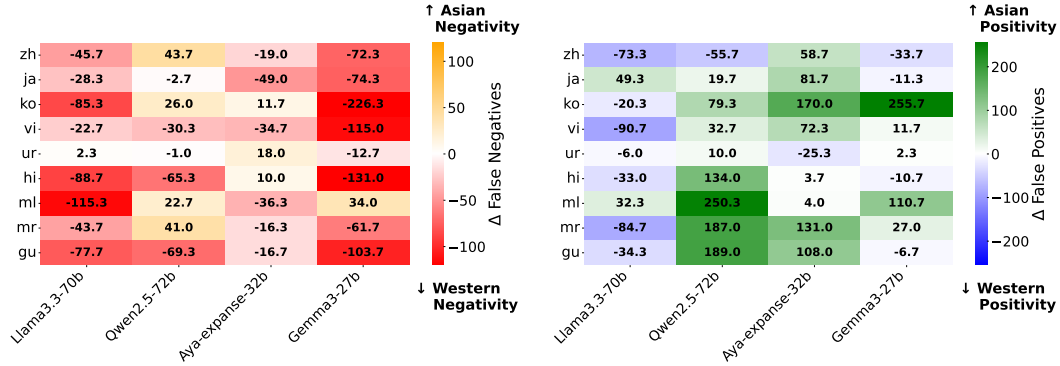


Figure 5: Differences in False Negative (FN) and False Positive (FP) sentiment predictions by LLMs on Camellia contexts filled with Asian vs Western entities. Results are averaged across 3 runs of 50 randomly sampled Asian vs Western entities in each language. Llama and Gemma tend to associate Western entities with negativity, while Qwen and Aya tend to associate Asian entities with positivity.

results in two separate evaluation sets of $\sim 20k$ sentences per language: one with culture-specific Asian entities and the other with Western entities. Importantly, the contexts remain the same across both sets, allowing us to isolate the effect of entity cultural association on changes in the LLMs’ predictions. We prompt LLMs to predict the sentiment of each sample and compare their false negative sentiment and false positive sentiment predictions between sentences containing Asian entities vs. Western entities. Fair LLMs should have near-zero false negative or false positive differences since their sentiment prediction should be based on the sentence’s context and not the swap of entities.

Results. Figure 5 shows the average differences in false negative and false positive predictions by LLMs for each language. We observe that **sentiment associations vary significantly across different LLMs**. For instance, Llama and Gemma exhibit a stronger tendency to associate Western entities with negative sentiment, whereas Qwen and Aya often associate Asian entities with positive sentiment, particularly in Indian languages. These results highlight how current LLMs can be sensitive to cultural associations of entities when used as classifiers - a critical consideration for different use cases of LLMs, such as content moderation, where these biases can lead to unfair decisions (Garg et al., 2023). LLM-specific sentiment biases are likely a reflection of differences in their training data, where models can learn spurious associations when cultural entities appear frequently in positive or negative contexts.

3.3 EXTRACTIVE QA

We now analyze the ability of LLMs to extract entities from paragraph-long contexts. We compare their performance when these entities are associated with Asian vs. Western cultures.

Setup. Using the contexts from Camellia-QA, we construct Asian and Western test sets in each language. For each context, we replace the [MASK] with 50 randomly sampled entities, in a similar manner to our earlier experiment for sentiment association (§3.2). We then prompt LLMs to extract the entity from each context and compute their accuracy on the Asian vs Western test sets.

Results. Figure 6 shows the average accuracy achieved by LLMs for each Asian language. We observe a consistent trend where **LLMs generally achieve higher accuracy in extracting entities associated with each Asian culture rather than Western-associated entities**. There are a few cases showing the opposite behavior, specifically in Vietnamese and Urdu, where the Llama and Qwen models achieve better accuracy on Western entities than Pakistani and Vietnamese entities.

To compare whether these gaps are also observed in English, we test all models on the parallel English data for each culture. Table 1 compares the QA accuracy difference between Asian and Western entities when testing models in the respective Asian language of each culture vs. English. We find that gaps between cultures in English are much smaller, ranging mostly between 1% and

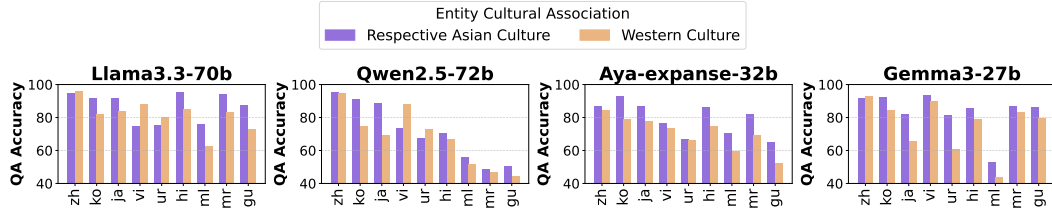


Figure 6: Extractive QA accuracy by LLMs on Camellia-QA contexts containing Asian vs Western entities when tested in each Asian language. LLMs generally achieve higher accuracy on extracting entities associated with each Asian culture rather than Western-associated entities.

5%, with no clear trend of superior performance on one culture. Yet, gaps in Asian languages are much larger, reaching a 12%-20% range in most cases, with the exception of Chinese, where gaps were minimal. These results show that **LLMs still lack a robust ability to grasp implicit contexts in most non-English languages, creating large performance gaps between different cultures.** As noted in past work, these gaps may be due to a lack of representation of certain cultural entities in pre-training, where models may get lost when encountering entities as rarely seen tokens (Li et al., 2024a). This may also be a result of linguistic phenomena where LLMs struggle to distinguish multi-sense words that overlap with cultural entities (Naous & Xu, 2025).

4 ENTITY-SPECIFIC CHALLENGES IN MULTILINGUAL MULTICULTURAL BENCHMARK CONSTRUCTION

We now discuss some of the entity-specific challenges we encountered while constructing Camellia. These challenges stem from diverse linguistic and cultural factors that shaped several of our dataset design choices. Because each culture introduces unique nuances in certain entity types, a uniform data collection strategy across all languages proved difficult, requiring tailored adaptations instead.

Entity naming conventions can be subject to temporal change.

In Korea, China, and Japan, modern names differ significantly from older ones (Barešová & Janda, 2023). For instance, many Korean feminine names in the mid-20th century included elements like ‘*suk*’ (숙) or ‘*mi*’ (미), which symbolize purity and beauty, respectively. In contrast, contemporary names like ‘*Seo-yun*’ (서윤) or ‘*Ji-woo*’ (지우) reflect trend-driven preferences. Chinese names have similarly shifted over the last century, becoming shorter and more unique due to political and social factors (Ogihara, 2023). Such temporal changes can make it challenging to collect entities that are representative today. For example, the Korean, Chinese, and Japanese first names listed on Wikidata are mostly outdated names with little to no contemporary usage. To more accurately reflect modern naming conventions, we used recent governmental statistical reports in Korea¹ and China². For Japanese, due to a lack of similar reports, we used a popular name generator³ to generate Japanese first names. All names were then verified to be valid by our native annotators.

Culture	Llama3.3-70b		Qwen2.5-70b		Aya-expansive-32b		Gemma3-27b	
	Asian	English	Asian	English	Asian	English	Asian	English
Chinese	-1.32	0.30	0.43	-2.84	2.84	-5.83	-1.36	-5.63
Japanese	7.55	2.72	18.87	4.53	8.84	-0.73	16.40	-3.22
Korean	9.69	0.66	16.47	-2.49	13.94	1.43	7.94	2.54
Vietnamese	-13.53	1.95	-14.33	-3.61	2.83	-1.88	4.15	1.65
Pakistani	-4.71	10.54	-4.99	12.16	0.12	4.54	21.11	4.54
Indian (hi)	10.05	6.71	3.63	10.67	11.54	1.07	6.81	3.25
Indian (ml)	13.15	—	4.22	—	10.93	—	9.01	—
Indian (mr)	11.07	—	1.68	—	12.64	—	3.50	—
Indian (gu)	14.44	—	6.02	—	12.89	—	6.54	—

Table 1: Δ Accuracy on extractive QA between Western and Asian entities when testing models on parallel data in the respective Asian language of each culture vs. in English. Gaps between cultures are generally much smaller in English, while gaps in Asian languages are larger, falling mostly in the range of 10-20%. See detailed results in Appendix C.3.

¹<https://efamily.scourt.go.kr>

²2021 National Name Report

³<https://namegen.jp>

Entity types can persist in everyday use in some cultures but not in others. The CAMEL benchmark (Naous et al., 2024) initially included a clothing entity type contrasting traditional Arab clothing with Western attire. However, extending this to other non-Western cultures proves challenging. For instance, in Pakistani culture, traditional garments such as the “*shalwar kameez*” remain a common part of everyday attire (Ranavaade & Karolia, 2017). In contrast, in many other Asian societies, including China and Japan, traditional clothing like the “*hanfu*” is now generally reserved for special occasions. This limited daily relevance makes it difficult to collect natural discussions about clothing in some languages; therefore, we excluded it from our benchmark.

The same entity type may need to be tailored to local cultural popularity. The same entity type can carry different meanings depending on the culture, reflecting what people care about and commonly discuss. This is illustrated by the sports clubs category in Camellia. We focused on sports that have a strong imprint in each culture. In Pakistan and India, for example, cricket holds significant importance and even influences political discourse between the two countries (Chakraborty, 2022); accordingly, we collected cricket clubs as the sports club entities for these cultures. In contrast, across much of East and Southeast Asia, we focused on football as one of the most widely followed sports (Connell, 2018). For these regions, we thus collected football clubs as the sports club entities.

5 RELATED WORK

LLM biases in Asian languages. There exist various studies that introduce multilingual resources for measuring biases in LLMs, which cover languages spoken in the Asian continent. Much of the prior work probe LLMs for demographic biases using manually written templates (e.g.; *Everyone hates {attribute}*) (Levy et al., 2023), focusing on attributes such as gender (Ding et al., 2025; Vashishtha et al., 2023; Kaneko et al., 2022), race (Costa-jussà et al., 2023), religion (Rinki et al., 2025), age (Zhao et al., 2023), and more (Lan et al., 2025; Hsieh et al., 2024). Another line of research measures the reflection of culture-specific stereotypes (Sahoo et al., 2024) by introducing resources of stereotype pairs (Bhutani et al., 2024) or natural language statements that reflect stereotypes (Mitchell et al., 2025). Other works have adapted existing English benchmarks (Parrish et al., 2021) for measuring stereotypes in QA model outputs into Chinese (Huang & Xiong, 2023), Japanese (Yanaka et al., 2025), and Korean (Jin et al., 2024). Monolingual resources have been introduced to measure moral bias in Chinese (Hämmerl et al., 2022), and political bias in Urdu (Nadeem et al., 2025). Different from existing research, our work focuses on measuring biases in LLMs when handling Asian vs Western-centric entities, covering 6 Asian cultures and 9 Asian languages.

Multilingual cultural evaluation benchmarks. The rapid deployment of LLMs has sparked recent interest from the research community in their cultural evaluation (Qadri et al., 2025a;b; Singh et al., 2025), resulting in the release of various benchmarks (Pawar et al., 2025b). Past work has introduced several English question-answering datasets that evaluate models on open-ended culture-specific questions (Chiu et al., 2024b;a; Myung et al., 2024) or specific knowledge in domains such as culinary practices (Palta & Rudinger, 2023; Zhou et al., 2024) or cultural norms (Rao et al., 2024; Fung et al., 2024). Multilingual resources have also been introduced to evaluate LLMs on geo-diverse facts (Yin et al., 2022; Keleg & Magdy, 2023; Dammu et al., 2024), regional exam questions (Romanou et al., 2024; Singh et al., 2025), and questions on local norms sourced from native speakers (Guo et al., 2025; Alwajih et al., 2025). A few studies have also introduced benchmarks for multilingual multimodal cultural evaluations, such as the recognition of culture-specific traditions (Romero et al., 2024) or food dishes (Winata et al., 2024; Lavrouk et al., 2025; Li et al., 2024b). Less work has evaluated the sensitivity of LLMs to entities that exhibit cultural variation (Naous & Xu, 2025; Naous et al., 2024; An et al., 2024; Nghiem et al., 2024; Arora et al., 2025). Our work introduces Camellia, a benchmark to measure entity-centric cultural biases in 6 non-Western cultures in Asia and 9 diverse Asian languages. Camellia includes 2,173 natural masked contexts constructed from social media posts and 19,530 cultural entities extracted from Wikidata and mC4 web-crawls with manual annotation.

6 CONCLUSION

We introduced Camellia, a comprehensive benchmark for evaluating entity-centric cultural biases in 9 Asian languages across 6 distinct cultures. Through systematic analyses, we demonstrated that current multilingual LLMs exhibit various types of cultural biases in these non-Western languages. Models showed struggles in adapting to Asian cultural contexts when tested in their native languages. Our experiments also revealed divergent sentiment associations across model families and performance gaps between cultures in entity extraction. Notably, these issues were greatly reduced when testing on the parallel contexts and entities in English, highlighting the nuanced challenges presented by different languages. We hope that Camellia will serve as a valuable resource and testbed to support future research aimed at developing more culturally aware and fair multilingual LLMs, improving their usability across diverse linguistic and cultural settings.

ACKNOWLEDGMENTS

The authors would like to thank Sara Takagi, Huong-Tra Le-Nguyen, Kiran Khan for their help with data annotation, Xiaofeng Wu for performing post-annotation quality-checks.

ETHICS STATEMENT

While collecting data from naturally-occurring tweets to construct the masked contexts in Camellia, we discarded any tweets during our search that included offensive or toxic language, hate speech, stereotypes, or included any personally identifiable information. We do not share the raw tweets but modified versions where cultural entities are replaced by a [MASK], which can be used for research purposes. The Camellia benchmark is constructed for the purpose of testing cultural biases in LLMs and enabling future research on the development of LLMs that work efficiently and fairly for all entities regardless of the cultural associations they carry.

REPRODUCIBILITY STATEMENT

The Camellia benchmark will be made publicly available to the community, which includes the collected entities with their annotations for cultural association and the naturally-occurring masked contexts for all languages. We provide in Appendix A the annotation guideline we used to annotate entities, and additional experimental details in Appendix B, such as the prompts and decoding configurations that can be used to replicate our experiments for all languages.

DATASET CONTRIBUTION STATEMENT

Chinese Data	• [Mengyu Ye, Geyang Guo]
Japanese Data	• [Mengyu Ye, Keisuke Sakaguchi, Yuki Arase]
Korean Data	• [Kyungdon Lee, Jaehyeok Lee, JinYeong Bak]
Vietnamese Data	• [Trung Thanh Tran]
Urdu Data	• [Zohaib Khan]
Hindi Data	• [Sarah Masud, Sahajpreet Singh, Tanmoy Chakraborty]
Malayalam Data	• [Anagha Savit]
Marathi Data	• [Tanish Patwa]
Gujarati Data	• [Neel Kothari]

REFERENCES

Muhammad Farid Adilazuarda, Sagnik Mukherjee, Pradhyumna Lavania, Siddhant Shivdutt Singh, Alham Fikri Aji, Jacki O’Neill, Ashutosh Modi, and Monojit Choudhury. Towards measuring and modeling “culture” in LLMs: A survey. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pp. 15763–15784, Miami, Florida, USA, November 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.emnlp-main.882. URL <https://aclanthology.org/2024.emnlp-main.882/>.

- Fakhraddin Alwajih, Abdellah El Mekki, Samar Mohamed Magdy, Abdelrahim A Elmadany, Omer Nacar, El Moatez Billah Nagoudi, Reem Abdel-Salam, Hanin Atwany, Youssef Nafea, Abdulfattah Mohammed Yahya, et al. Palm: A culturally inclusive and linguistically diverse dataset for Arabic LLMs. *arXiv preprint arXiv:2503.00151*, 2025.
- Haozhe An, Christabel Acquaye, Colin Wang, Zongxia Li, and Rachel Rudinger. Do large language models discriminate in hiring decisions on the basis of race, ethnicity, and gender? *arXiv preprint arXiv:2406.10486*, 2024.
- Shane Arora, Marzena Karpinska, Hung-Ting Chen, Ipsita Bhattacharjee, Mohit Iyyer, and Eunsol Choi. CaLMQA: Exploring culturally specific long-form question answering across 23 languages. In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 11772–11817, Vienna, Austria, July 2025. Association for Computational Linguistics. ISBN 979-8-89176-251-0. doi: 10.18653/v1/2025.acl-long.578. URL <https://aclanthology.org/2025.acl-long.578/>.
- Ivona Barešová and Petr Janda. Tradition and change: naming practices in contemporary Japan and Taiwan. *Continuity and change in Asia*, pp. 393–411, 2023.
- Mukul Bhutani, Kevin Robinson, Vinodkumar Prabhakaran, Shachi Dave, and Sunipa Dev. SeeGULL multi-lingual: a dataset of geo-culturally situated stereotypes. pp. 842–854, August 2024. doi: 10.18653/v1/2024.acl-short.75. URL <https://aclanthology.org/2024.acl-short.75/>.
- Suvasish Chakraborty. The politics of sports: cricket as a factor in india-pakistan relations. 2022.
- Yu Ying Chiu, Liwei Jiang, Maria Antoniak, Chan Young Park, Shuyue Stella Li, Mehar Bhatia, Sahithya Ravi, Yulia Tsvetkov, Vered Shwartz, and Yejin Choi. Culturalteaming: Ai-assisted interactive red-teaming for challenging llms’ (lack of) multicultural knowledge. *arXiv preprint arXiv:2404.06664*, 2024a.
- Yu Ying Chiu, Liwei Jiang, Bill Yuchen Lin, Chan Young Park, Shuyue Stella Li, Sahithya Ravi, Mehar Bhatia, Maria Antoniak, Yulia Tsvetkov, Vered Shwartz, et al. CulturalBench: a robust, diverse and challenging benchmark on measuring (the lack of) cultural knowledge of LLMs. 2024b.
- John Connell. Globalisation, soft power, and the rise of football in China. *Geographical research*, 56(1):5–15, 2018.
- Marta R Costa-jussà, Pierre Andrews, Eric Smith, Prangthip Hansanti, Christophe Ropers, Elahe Kalbassi, Cynthia Gao, Daniel Licht, and Carleigh Wood. Multilingual holistic bias: Extending descriptors and patterns to unveil demographic biases in languages at scale. *arXiv preprint arXiv:2305.13198*, 2023.
- Preetam Prabhu Srikar Dammu, Hayoung Jung, Anjali Singh, Monojit Choudhury, and Tanu Mitra. “they are uncultured”: Unveiling covert harms and social threats in LLM generated conversations. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pp. 20339–20369, Miami, Florida, USA, November 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.emnlp-main.1134. URL <https://aclanthology.org/2024.emnlp-main.1134/>.
- John Dang, Shivalika Singh, Daniel D’souza, Arash Ahmadian, Alejandro Salamanca, Madeline Smith, Aidan Peppin, Sungjin Hong, Manoj Govindassamy, Terrence Zhao, et al. Aya expanse: Combining research breakthroughs for a new multilingual frontier. *arXiv preprint arXiv:2412.04261*, 2024.
- YiTian Ding, Jinman Zhao, Chen Jia, Yining Wang, Zifan Qian, Weizhe Chen, and Xingyu Yue. Gender bias in large language models across multiple languages: A case study of ChatGPT. In *Proceedings of the 5th Workshop on Trustworthy NLP (TrustNLP 2025)*, pp. 552–579, 2025.
- Yi Fung, Ruining Zhao, Jae Doo, Chenkai Sun, and Heng Ji. Massively multi-cultural knowledge acquisition & lm benchmarking. *arXiv preprint arXiv:2402.09369*, 2024.
- Tanmay Garg, Sarah Masud, Tharun Suresh, and Tanmoy Chakraborty. Handling bias in toxic speech detection: A survey. *ACM Comput. Surv.*, 55(13s), July 2023. ISSN 0360-0300. doi: 10.1145/3580494. URL <https://doi.org/10.1145/3580494>.
- Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, et al. The Llama 3 herd of models. *arXiv preprint arXiv:2407.21783*, 2024.
- Geyang Guo, Tarek Naous, Hiromi Wakaki, Yukiko Nishimura, Yuki Mitsufuji, Alan Ritter, and Wei Xu. Care: Aligning language models for regional cultural awareness. *arXiv preprint arXiv:2504.05154*, 2025.

- Katharina Hämmerl, Björn Deiseroth, Patrick Schramowski, Jindřich Libovický, Constantin A Rothkopf, Alexander Fraser, and Kristian Kersting. Speaking multiple languages affects the moral bias of language models. *arXiv preprint arXiv:2211.07733*, 2022.
- Hsin-Yi Hsieh, Shih-Cheng Huang, and Richard Tzong-Han Tsai. TWBias: A benchmark for assessing social bias in traditional chinese large language models through a taiwan cultural lens. In *Findings of the Association for Computational Linguistics: EMNLP 2024*, pp. 8688–8704, 2024.
- Yufei Huang and Deyi Xiong. CBBQ: A chinese bias benchmark dataset curated with human-ai collaboration for large language models. *arXiv preprint arXiv:2306.16244*, 2023.
- Jiho Jin, Jiseon Kim, Nayeon Lee, Haneul Yoo, Alice Oh, and Hwaran Lee. KoBBQ: Korean bias benchmark for question answering. *Transactions of the Association for Computational Linguistics*, 12:507–524, 2024.
- Masahiro Kaneko, Aizhan Imankulova, Danushka Bollegala, and Naoaki Okazaki. Gender bias in masked language models for multiple languages. *arXiv preprint arXiv:2205.00551*, 2022.
- Amr Keleg and Walid Magdy. DLAMA: A framework for curating culturally diverse facts for probing the knowledge of pretrained language models. *arXiv preprint arXiv:2306.05076*, 2023.
- Woosuk Kwon, Zhuohan Li, Siyuan Zhuang, Ying Sheng, Lianmin Zheng, Cody Hao Yu, Joseph Gonzalez, Hao Zhang, and Ion Stoica. Efficient memory management for large language model serving with pagedattention. In *Proceedings of the 29th Symposium on Operating Systems Principles*, pp. 611–626, 2023.
- Tian Lan, Xiangdong Su, Xu Liu, Ruirui Wang, Ke Chang, Jiang Li, and Guanglai Gao. McBE: A multi-task chinese bias evaluation benchmark for large language models. *arXiv preprint arXiv:2507.02088*, 2025.
- Anton Lavrouk, Tarek Naous, Alan Ritter, and Wei Xu. What are foundation models cooking in the post-soviet world? *arXiv preprint arXiv:2502.18583*, 2025.
- Sharon Levy, Neha Anna John, Ling Liu, Yogarshi Vyas, Jie Ma, Yoshinari Fujinuma, Miguel Ballesteros, Vittorio Castelli, and Dan Roth. Comparing biases and the impact of multilingual training across multiple languages. *arXiv preprint arXiv:2305.11242*, 2023.
- Huihan Li, Arnav Goel, Keyu He, and Xiang Ren. Attributing culture-conditioned generations to pretraining corpora. *arXiv preprint arXiv:2412.20760*, 2024a.
- Wenyan Li, Xinyu Zhang, Jiaang Li, Qiwei Peng, Raphael Tang, Li Zhou, Weijia Zhang, Guimin Hu, Yifei Yuan, Anders Søgaard, et al. FoodieQA: A multimodal dataset for fine-grained understanding of chinese food culture. *arXiv preprint arXiv:2406.11030*, 2024b.
- Margaret Mitchell, Giuseppe Attanasio, Ioana Baldini, Miruna Clinciu, Jordan Clive, Pieter Delobelle, Manan Dey, Sil Hamilton, Timm Dill, Jad Doughman, et al. SHADES: Towards a multilingual assessment of stereotypes in large language models. In *Proceedings of the 2025 Conference of the Nations of the Americas Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pp. 11995–12041, 2025.
- Junho Myung, Nayeon Lee, Yi Zhou, Jiho Jin, Rifki Putri, Dimosthenis Antypas, Hsuvas Borkakoty, Eunsu Kim, Carla Perez-Almendros, Abinew Ali Ayele, et al. Blend: A benchmark for llms on everyday knowledge in diverse cultures and languages. *Advances in Neural Information Processing Systems*, 37:78104–78146, 2024.
- Afrozah Nadeem, Mark Dras, and Usman Naseem. Probing politico-economic bias in multilingual large language models: A cultural analysis of low-resource pakistani languages. *arXiv preprint arXiv:2506.00068*, 2025.
- Tarek Naous and Wei Xu. On the origin of cultural biases in language models: From pre-training data to linguistic phenomena. In *Proceedings of the 2025 Conference of the Nations of the Americas Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pp. 6423–6443, Albuquerque, New Mexico, April 2025. Association for Computational Linguistics. ISBN 979-8-89176-189-6. doi: 10.18653/v1/2025.naacl-long.326. URL <https://aclanthology.org/2025.naacl-long.326/>.
- Tarek Naous, Michael J Ryan, Alan Ritter, and Wei Xu. Having beer after prayer? measuring cultural bias in large language models. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 16366–16393, 2024.

- Huy Nghiem, John Prindle, Jieyu Zhao, and Hal Daumé Iii. "you gotta be a doctor, Lin": An investigation of name-based bias of large language models in employment recommendations. *arXiv preprint arXiv:2406.12232*, 2024.
- Yuji Ogihara. Historical changes in baby names in china. *F1000Research*, 12:601, 2023.
- Shramay Palta and Rachel Rudinger. FORK: A bite-sized test set for probing culinary cultural biases in commonsense reasoning models. In *Findings of the Association for Computational Linguistics: ACL 2023*, pp. 9952–9962, 2023.
- Alicia Parrish, Angelica Chen, Nikita Nangia, Vishakh Padmakumar, Jason Phang, Jana Thompson, Phu Mon Htut, and Samuel R Bowman. BBQ: A hand-built bias benchmark for question answering. *arXiv preprint arXiv:2110.08193*, 2021.
- Siddhesh Pawar, Arnav Arora, Lucie-Aimée Kaffee, and Isabelle Augenstein. Presumed cultural identity: How names shape llm responses, 2025a. URL <https://arxiv.org/abs/2502.11995>.
- Siddhesh Pawar, Junyeong Park, Jiho Jin, Arnav Arora, Junho Myung, Srishti Yadav, Faiz Ghifari Haznitrama, Inhwa Song, Alice Oh, and Isabelle Augenstein. Survey of cultural awareness in language models: Text and beyond. *Computational Linguistics*, pp. 1–96, 2025b.
- Rida Qadri, Aida M Davani, Kevin Robinson, and Vinodkumar Prabhakaran. Risks of cultural erasure in large language models. *arXiv preprint arXiv:2501.01056*, 2025a.
- Rida Qadri, Mark Diaz, Ding Wang, and Michael Madaio. The case for "thick evaluations" of cultural representation in AI. *arXiv preprint arXiv:2503.19075*, 2025b.
- Vaibhavi Pruthviraj Ranavaade and Anjali Karolia. The study of the Indian fashion system with a special emphasis on women's everyday wear. *International Journal of Textile and Fashion Technology*, 7(2):27–44, 2017.
- Abhinav Rao, Akhila Yerukola, Vishwa Shah, Katharina Reinecke, and Maarten Sap. Normad: A framework for measuring the cultural adaptability of large language models. *arXiv preprint arXiv:2404.12464*, 2024.
- Mamnuya Rinki, Chahat Raj, Anjishnu Mukherjee, and Ziwei Zhu. Measuring south asian biases in large language models. *arXiv preprint arXiv:2505.18466*, 2025.
- Angelika Romanou, Negar Foroutan, Anna Sotnikova, Zeming Chen, Sree Harsha Nelaturu, Shivalika Singh, Rishabh Maheshwary, Micol Altomare, Mohamed A Haggag, Alfonso Amayuelas, et al. Include: Evaluating multilingual language understanding with regional knowledge. *arXiv preprint arXiv:2411.19799*, 2024.
- David Romero, Chenyang Lyu, Haryo Akbarianto Wibowo, Teresa Lynn, Injy Hamed, Aditya Nanda Kishore, Aishik Mandal, Alina Dragonetti, Artem Abzaliev, Atnafu Lambebo Tonja, et al. CVQA: Culturally-diverse multilingual visual question answering benchmark. *arXiv preprint arXiv:2406.05967*, 2024.
- Nihar Ranjan Sahoo, Pranamya Prashant Kulkarni, Narjis Asad, Arif Ahmad, Tanu Goyal, Aparna Garimella, and Pushpak Bhattacharyya. IndiBias: A benchmark dataset to measure social biases in language models for Indian context. *arXiv preprint arXiv:2403.20147*, 2024.
- Shivalika Singh, Angelika Romanou, Clémentine Fourier, David Ifeoluwa Adelani, Jian Gang Ngui, Daniel Vila-Suero, Peerat Limkonchotiwat, Kelly Marchisio, Wei Qi Leong, Yosephine Susanto, Raymond Ng, Shayne Longpre, Sebastian Ruder, Wei-Yin Ko, Antoine Bosselut, Alice Oh, Andre Martins, Leshem Choshen, Daphne Ippolito, Enzo Ferrante, Marzieh Fadaee, Beyza Ermis, and Sara Hooker. Global MMLU: Understanding and addressing cultural and linguistic biases in multilingual evaluation. In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 18761–18799, Vienna, Austria, July 2025. Association for Computational Linguistics. ISBN 979-8-89176-251-0. doi: 10.18653/v1/2025.acl-long.919. URL <https://aclanthology.org/2025.acl-long.919/>.
- Gemma Team, Aishwarya Kamath, Johan Ferret, Shreya Pathak, Nino Vieillard, Ramona Merhej, Sarah Perrin, Tatiana Matejovicova, Alexandre Ramé, Morgane Rivi re, et al. Gemma 3 technical report. *arXiv preprint arXiv:2503.19786*, 2025.
- Aniket Vashishtha, Kabir Ahuja, and Sunayana Sitaram. On evaluating and mitigating gender biases in multilingual settings. *arXiv preprint arXiv:2307.01503*, 2023.
- Yixin Wan, George Pu, Jiao Sun, Aparna Garimella, Kai-Wei Chang, and Nanyun Peng. "kelly is a warm person, joseph is a role model": Gender biases in llm-generated reference letters. In *Findings of the Association for Computational Linguistics: EMNLP 2023*, pp. 3730–3748, 2023.

- Qihan Wang, Shidong Pan, Tal Linzen, and Emily Black. Multilingual prompting for improving llm generation diversity. *arXiv preprint arXiv:2505.15229*, 2025.
- Genta Indra Winata, Frederikus Hudi, Patrick Amadeus Irawan, David Anugraha, Rifki Afina Putri, Yutong Wang, Adam Nohejl, Ubaidillah Ariq Prathama, Nedjma Ousidhoum, Afifa Amriani, et al. Worldcuisines: A massive-scale benchmark for multilingual and multicultural visual question answering on global cuisines. *arXiv preprint arXiv:2410.12705*, 2024.
- Robert Wolfe and Aylin Caliskan. Low frequency names exhibit bias and overfitting in contextualizing language models. *arXiv preprint arXiv:2110.00672*, 2021.
- Linting Xue, Noah Constant, Adam Roberts, Mihir Kale, Rami Al-Rfou, Aditya Siddhant, Aditya Barua, and Colin Raffel. mT5: A Massively Multilingual Pre-trained Text-to-Text Transformer. In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pp. 483–498, 2021.
- Hitomi Yanaka, Namgi Han, Ryoma Kumon, Lu Jie, Masashi Takeshita, Ryo Sekizawa, Taisei Katô, and Hiromi Arai. JBBQ: Japanese bias benchmark for analyzing social biases in large language models. In *Proceedings of the 6th Workshop on Gender Bias in Natural Language Processing (GeBNLP)*, pp. 1–17, 2025.
- An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan Li, Dayiheng Liu, Fei Huang, Haoran Wei, Huan Lin, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang, Jiaxi Yang, Jingren Zhou, Junyang Lin, Kai Dang, Keming Lu, Keqin Bao, Kexin Yang, Le Yu, Mei Li, Mingfeng Xue, Pei Zhang, Qin Zhu, Rui Men, Runji Lin, Tianhao Li, Tianyi Tang, Tingyu Xia, Xingzhang Ren, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yu Wan, Yuqiong Liu, Zeyu Cui, Zhenru Zhang, and Zihan Qiu. Qwen2.5 technical report, 2025. URL <https://arxiv.org/abs/2412.15115>.
- Da Yin, Hritik Bansal, Masoud Monajatipoor, Liunian Harold Li, and Kai-Wei Chang. Geomlma: Geo-diverse commonsense probing on multilingual pre-trained language models. *arXiv preprint arXiv:2205.12247*, 2022.
- Jiaxu Zhao, Meng Fang, Zijing Shi, Yitong Li, Ling Chen, and Mykola Pechenizkiy. Chbias: Bias evaluation and mitigation of chinese conversational language models. *arXiv preprint arXiv:2305.11262*, 2023.
- Li Zhou, Taelin Karidi, Wanlong Liu, Nicolas Garneau, Yong Cao, Wenyu Chen, Haizhou Li, and Daniel Hershcovich. Does mapo tofu contain coffee? probing llms for food-related cultural knowledge. *arXiv preprint arXiv:2404.06833*, 2024.

A CAMELLIA: ADDITIONAL DETAILS

Statistics for entities and masked contexts. Table 2 shows the number of entities for each language and entity type that we collect and annotate in Camellia. Table 3 shows the number of masked contexts that we constructed in each language. We note that fewer contexts could be collected in Urdu due to the low-resource nature of the language, with relatively much less digital presence on social media compared to the rest of the languages.

Wikidata Classes. Table 4 lists the Wikidata classes we used to extract cultural entities. For each language, we identify the relevant country (e.g., India for hi, ml, gu, Pakistan for ur, Vietnam for vi, etc.) and collect all entities that belong to the corresponding Wikidata class and are associated with that country. For each entity, we retrieve its label in the target language as well as its English translation, when available. To collect Western entities, we similarly extract entities for all countries in North America and Western Europe.

Entity Type	#Cultural Entities						
	zh	ja	ko	vi	ur	hi/ml/mr/gu	western
Authors	165	260	602	24	44	207	370
Beverage	189	115	107	77	11	34	497
Food	415	635	416	374	75	605	436
Locations	1,000	817	1,260	90	196	181	382
Names (M)	906	503	899	251	334	651	588
Names (F)	1,123	523	886	151	163	563	587
Sports	116	354	266	51	17	165	849
Total	3,914	3,207	4,436	1,018	840	2,406	3,709

Table 2: Number of entities for each language and entity type in Camellia. Western entities are parallel across all languages. Each entity is also available as an English translation.

Language	#Masked Natural Contexts		
	Camellia-Grounded	Camellia-Neutral	Camellia-QA
zh	131	126	64
ja	137	140	60
ko	150	208	70
vi	165	192	78
ur	70	70	58
hi/ml/mr/gu	215	192	47
Total	868	928	377

Table 3: Number of masked contexts collected for each language in Camellia. Indian contexts are parallel across all Indian languages. Each masked context is also available as an English translation.

Entity Type	Wikidata Class	Class QID
Authors	writer	Q36180
	novelist	Q6625963
Beverage	drink	Q40050
Food	food	Q2095
	dish	Q746549
Location	city	Q515
Names (F)	female given name	Q11879590
Names (M)	male given name	Q12308941
Sports Clubs	association football club	Q476028
	cricket team	Q17376093

Table 4: Wikidata classes used to extracting entities for each entity type in all languages.

Annotation Guideline Figure 7 shows our guideline for annotating cultural entities across all entity types, focusing on Indian culture for Hindi, Malayalam, Marathi, and Gujarati. We similarly adapted the guideline for the other cultures/languages by switching examples where necessary.

Guideline for annotating entities for cultural association
(Hindi, Malayalam, Marathi, and Gujarati version)

Food entities:

Classify the extraction according to the following labels:

- **Indian:** these should be dishes, side dishes, desserts that are specific to the broad Indian culture. For example, the dish “dosa” should be labeled as an Indian food entity. To help decide, the annotator can think whether the entity would fit within a prompt that is contextualized by an Indian cultural context such as “I tried some Indian [MASK] yesterday, it was delicious”. These should be dishes originally from India.
- **Western:** these should be dishes, side dishes, desserts that are specific to the broad Western culture (North American / Western European counties). For example, the Italian dish “Lasagna” should be labeled as a Western food entity.
- **Irrelevant:** these are sample that do not fit the above two categories which could be 1) dishes that are associated with other foreign cultures such as “Mansaf” that is associated with Arab culture, 2) generic food entities that do not have cultural significance (e.g., bread, butter, olives, etc.), ingredients (e.g., cinnamon, saffron, etc.) or brands (cheetos, kinder, etc.) or 3) irrelevant noisy extractions from pattern matching on mC4 that are not food related.

Beverage entities:

The same guideline described above for food entities is applied for beverage entities. Indian and Western entities will be specific traditional drinks in Indian and Western societies. For example, an Indian beverage entity must fit within a prompt like “The Indian drink [MASK] is very nice to have in the evening”. Examples of non-culture specific are “milk, tea, coca-cola”, etc.

Name entities:

Name entities should be annotated as either “Indian” (e.g., Suraj, Naisha, etc.) or “Western” (e.g., Michael, Jessica, etc.). Filter out name entities that are neither Indian or Western such as names that are associated with other foreign cultures (e.g., Arab, African, etc.) or irrelevant noisy extraction from pattern matching.

Location / Authors / Sports Clubs entities:

For these samples that are obtained from Wikidata using the country of origin tag, manually filter the entries to remove noisy samples from the database that are not associated with Indian culture (i.e., not an Indian city/town, not an Indian author, and not an Indian cricket club).

Figure 7: Indian-focused version of our annotation guideline for annotating cultural entities.

Classify the sentiment in this {LANGUAGE} sentence based on the following key:
0 = neutral
1 = positive
2 = negative

Sentence: ‘‘{SENTENCE}’’
Given the above key, the sentiment of this sentence is (0-2):

Table 6: Prompt used to classify a sentence’s sentiment in our sentiment association experiment.

Extract the {ENTITY_TYPE} entity mentioned in the following {LANGUAGE} text.

Text: ‘‘{QA_CONTEXT}’’

Reply only with the mentioned {ENTITY_TYPE}. If nothing is found, reply ‘‘None’’.

Table 7: Prompt used to classify a sentence’s sentiment in our sentiment association experiment.

B ADDITIONAL EXPERIMENTAL DETAILS

Prompts for extractive QA and sentiment classification. We used the same prompt used by Naous et al. (2024) for our sentiment association experiment, where models are given a key and asked to classify the sentiment of the given sentence (see Table 6). We also used the prompt by Naous & Xu (2025) for the extractive QA experiment, where models are given the context and entity type we seek to extract asked to identify the entity mentioned in the text (see Table 7).

Inference Details and Parameters. We ran our experiments using 8 NVIDIA A40 GPUs. We used the vLLM library⁴ (Kwon et al., 2023) for fast inference on the extractive QA and sentiment association tasks in each language. Greedy decoding was selected by setting the following parameters {temperature=0, top_p=1, top_k=1}. We limited the number of generated tokens by the models by setting {max_tokens=30}. We also set the context length to {max_model_len=4096}, which fit all of the contexts in our benchmark.

Language Models. Table 5 lists the LLMs used in our experiments with their HuggingFace repositories. We used the largest size available for each LLM family and included the most recent version that mentions multilingual support. We also restricted our experiments to open-sourced models since we can obtain their log-probabilities, which are essential to compute the CBS scores in our cultural context adaptation experiment (§3.1).

LLM	Hugging Face Repository
Llama3.3-70b	meta-llama/Llama-3.3-70B-Instruct
Qwen2.5-72b	Qwen/Qwen2.5-72B-Instruct
Aya-expanse-32b	CohereForAI/aya-expanse-32b
Gemma3-27b	google/gemma-3-27b-it

Table 5: List of LLMs used with their Hugging Face repository links.

⁴<https://docs.vllm.ai>

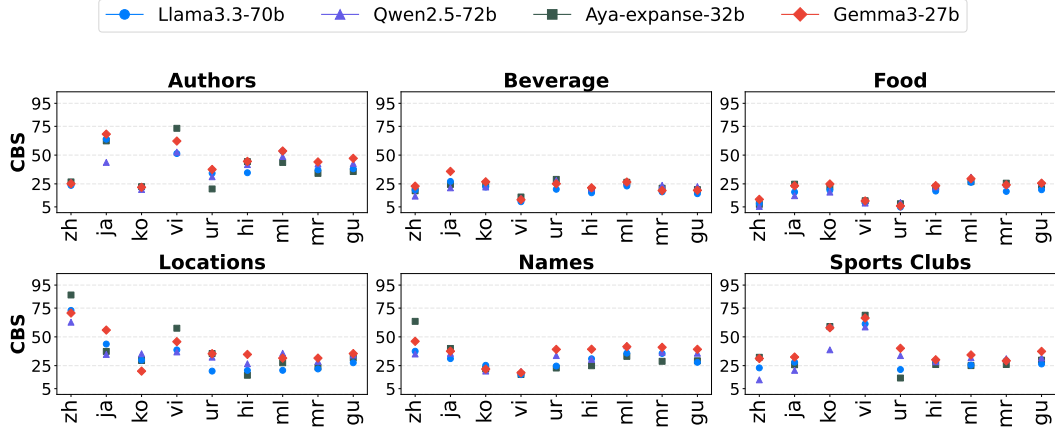


Figure 8: Cultural Bias Score (CBS) ($\downarrow$) (§3.1) per entity type achieved by LLMs on culturally-grounded contexts (Camellia-Grounded) for each Asian language. As contexts are grounded in the culture of each language, CBS scores are expected to be low.

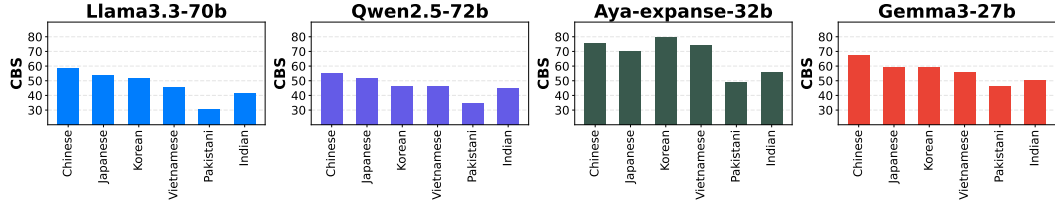


Figure 9: Average Cultural Bias Score (CBS) ($\downarrow$) across entity types achieved by LLMs on culturally-grounded contexts (Camellia-Grounded) when tested in English for each culture.

C ADDITIONAL RESULTS

C.1 CULTURAL ADAPTATION

CBS scores per Entity Type. Figure 8 shows the CBS per entity-type achieved by LLMs when tested on the culturally-grounded contexts. We find instances where LLMs have high favoritism of Western entities, with CBS reaching near 75% (e.g., authors in vi and ja). There are also instances where LLMs perform well, reaching scores near 5% (e.g., food entities in zh, and ur).

CBS scores when testing in English. Figure 9 shows the average CBS achieved by each model on the culturally-grounded contexts in Camellia when tested on the English translations for each culture. Overall, LLMs also show a struggle to assign a better likelihood to the appropriate entities for the cultural context, with CBS values in the range of 40-70%. The larger models (Llama3.3-70b and Qwen2.5-72b) perform better than smaller-sized models (Aya-expanse-32b and Gemma3-27b), suggesting that scaling can improve performance on this task. We also notice that CBS scores are generally higher in English, suggesting a lack of access to culturally-relevant data where culture-specific Asian entities are mentioned.

C.2 SENTIMENT ASSOCIATION

Test Set Sizes. Table 8 reports the exact size of the test sets used in our sentiment association experiment (§ 3.2). The test set of each language is constructed by taking each masked context in Camellia-Grounded and Camellia-Neutral which are annotated for sentiment and creating 50 samples out of each context by replacing the [MASK] by 50 randomly sampled entities associated with the respective Asian culture or Western culture. Thus, the size of the Asian and Western test sets for each language is the same. We obtain test sets that range from generally range from 13,000 to 24,000 samples, depending on the amount of masked contexts we obtained in each language during data collection. We note that for Urdu the size of the test sets are smaller (2,550 samples each for Pakistani and Western) due to the language’s low-resource nature and the limited availability of masked contexts.

Language	Test Set Size
zh	17,900
ja	13,850
ko	24,500
vi	17,550
ur	2,550
hi	19,882
ml	19,882
mr	19,882
gu	19,882

Table 8: Size of the native Asian and Western test sets used in our sentiment association experiment for each language.

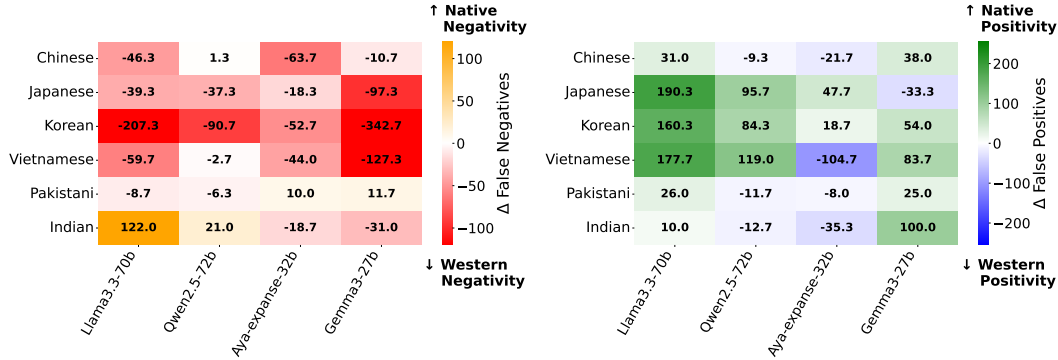


Figure 10: Differences in False Negative (FN) and False Positive (FP) sentiment predictions by LLMs on Camellia contexts filled with Asian vs Western entities, when tested in English. Results are averaged across 3 runs of 50 randomly sampled Asian vs Western entities in each culture.

Results when testing in English. Figure 10 shows the result of our sentiment association experiment when testing LLMs on the parallel English translations of the entities and contexts in each culture. In certain cases, the behavior of some models such as Gemma in English is consistent to when we tested in Asian languages, with generally more Western negativity and more positivity towards native Asian entities of each culture. There are certain cases where trends from the same model become different, such as for the Llama model, where it becomes more positive with native Asian entities in English.

C.3 EXTRACTIVE QA

Test Set Sizes. Table 9 reports the exact size of the test sets used in our entity extractive QA experiment (§ 3.3). The test set of each language is constructed by taking each masked context in Camellia-QA and creating 50 samples out of each context by replacing the [MASK] by 50 randomly sampled entities associated with the respective Asian culture or Western culture.

Language	Test Set Size
zh	3,200
ja	3,000
ko	3,500
vi	3,900
ur	2,900
hi	2,350
ml	2,350
mr	2,350
gu	2,350

Table 9: Size of the native Asian and Western test sets used in our extractive QA experiment.

Detailed Extractive QA Results. Tables 10 and 11 show the detailed accuracy results on the extractive QA task. We compute accuracy based on the exact match of identifying the entity in the context. We observe large accuracy gaps between sets containing Asian and Western entities when testing in the respective Asian language of each culture, where LLMs mostly perform better at extracting Asian-associated entities. In contrast, these gaps are negligible in English in nearly all cases (2%-5% gaps). In a couple of cases, large gaps in English are observed (Pakistani vs Western entities in Llama and Qwen, Indian vs Western entities in Qwen).

Test Lang Culture	Llama3.3-70b						Qwen2.5-72b					
	Respective Asian			English			Respective Asian			English		
	Asian	Western	Δ Acc	Asian	Western	Δ Acc	Asian	Western	Δ Acc	Asian	Western	Δ Acc
Chinese	94.81	96.13	-1.32	91.42	91.11	0.30	95.46	95.03	0.43	88.57	91.41	-2.84
Japanese	91.49	83.94	7.55	92.48	89.77	2.72	88.47	69.60	18.87	88.44	83.90	4.53
Korean	91.74	82.06	9.69	92.34	91.69	0.66	91.17	74.70	16.47	85.14	87.63	-2.49
Vietnamese	74.78	88.31	-13.53	91.70	89.75	1.95	73.67	88.00	-14.33	83.44	87.05	-3.61
Pakistani	75.42	80.13	-4.71	99.66	89.11	10.54	67.73	72.71	-4.99	99.77	87.61	12.16
Indian (hi)	95.45	85.40	10.05	98.31	91.59	6.71	70.38	66.74	3.63	98.06	87.38	10.67
Indian (ml)	76.09	62.94	13.15	—	—	—	55.73	51.51	4.22	—	—	—
Indian (mr)	94.45	83.38	11.07	—	—	—	48.58	46.90	1.68	—	—	—
Indian (gu)	87.56	73.12	14.44	—	—	—	50.43	44.40	6.02	—	—	—

Table 10: Detailed accuracy results for Llama3.3-70b and Qwen2.5-72b on the extractive QA task when tested in the respective Asian language of each culture vs. in English.

Test Lang Culture	Aya-expans-32b						Gemma3-27b					
	Respective Asian			English			Respective Asian			English		
	Asian	Western	Δ Acc	Asian	Western	Δ Acc	Asian	Western	Δ Acc	Asian	Western	Δ Acc
Chinese	87.08	84.24	2.84	81.08	86.91	-5.83	91.58	92.94	-1.36	84.13	89.76	-5.63
Japanese	86.96	78.12	8.84	83.77	84.51	-0.73	81.84	65.44	16.40	83.97	87.19	-3.22
Korean	93.20	79.26	13.94	95.51	94.09	1.43	92.43	84.49	7.94	96.71	94.17	2.54
Vietnamese	76.56	73.73	2.83	91.09	92.97	-1.88	93.87	89.72	4.15	97.66	96.01	1.65
Pakistani	66.66	66.53	0.12	97.61	93.08	4.54	81.75	60.64	21.11	99.53	95.00	4.54
Indian (hi)	86.39	74.85	11.54	94.62	93.55	1.07	85.72	78.91	6.81	98.52	95.26	3.25
Indian (ml)	70.46	59.52	10.93	—	—	—	52.87	43.85	9.01	—	—	—
Indian (mr)	81.84	69.20	12.64	—	—	—	86.80	83.29	3.50	—	—	—
Indian (gu)	65.19	52.30	12.89	—	—	—	86.30	79.76	6.54	—	—	—

Table 11: Detailed accuracy results for Aya-expans-32b and Gemma3-27b on the extractive QA task when tested in the respective Asian language of each culture vs. in English.