

Exposing Blindspots: Cultural Bias Evaluation in Generative Image Models

Huichan Seo^{1*} **Sieun Choi^{2*†}** **Minki Hong^{2*†}** **Yi Zhou¹**
Junseo Kim^{3†} **Lukman Ismaila⁴** **Naome Etori⁵** **Mehul Agarwal¹**
Zhixuan Liu¹ **Jihie Kim²** **Jean Oh^{1,6}**

¹Carnegie Mellon University, Pittsburgh, United States

²Dongguk University, Seoul, South Korea

³Delft University of Technology, Delft, Netherlands

⁴Johns Hopkins University, School of Medicine, Baltimore, United States

⁵University of Minnesota–Twin Cities, Minneapolis, United States

⁶Lavoro AI, Pittsburgh, United States

Corresponding author: Jean Oh (jeanoh@cmu.edu)

Abstract

Generative image models produce striking visuals yet often misrepresent culture. Prior work has probed cultural dimensions primarily in text-to-image (T2I) systems, leaving image-to-image (I2I) editors largely underexamined. We close this gap with a unified, reproducible evaluation spanning six countries, an 8-category/36-subcategory schema, and era-aware prompts, auditing both T2I generation and I2I editing under a standardized, reproducible protocol that yields comparable model-level diagnostics. Using open models with fixed configurations, we derive comparable diagnostics across countries, eras, and categories for both T2I and I2I. Our evaluation combines standard automatic measures, a culture-aware metric that integrates retrieval-augmented visual question answering (VQA) with curated knowledge, and expert human judgments collected on a web platform from country-native reviewers. To enable downstream analyses without re-running compute-intensive pipelines, we release the complete image corpus from both studies alongside prompts and settings. Our study reveals three recurring findings. First, under country-agnostic prompts, models default to U.S.-like, modern-leaning depictions and flatten cross-country distinctions, reducing separability between culturally distinct neighbors despite fixed schema and era controls. Second, iterative I2I editing erodes cultural fidelity even when conventional metrics remain flat or improve; by contrast, expert ratings and our culture-aware metric both register this degradation. Third, I2I models tend to apply superficial cues (palette shifts, generic props) rather than context- and era-consistent changes, frequently retaining source identity for non-U.S. targets and drifting toward non-photorealistic styles; attribute-addition trials further expose weak text rendering and brittle handling of fine, culture-specific details. Taken together, these results indicate that culture-sensitive edits remain unreliable in current systems. By standardizing prompts, settings, metrics, and human evaluation protocols—and releasing all images and configurations—we offer a reproducible, culture-centered pipeline for diagnosing and tracking progress in generative image research. Project page: <https://seochan99.github.io/ECB/>.

*Equal contribution.

†This paper is based on the work performed while the authors were visiting scholars at Carnegie Mellon University.

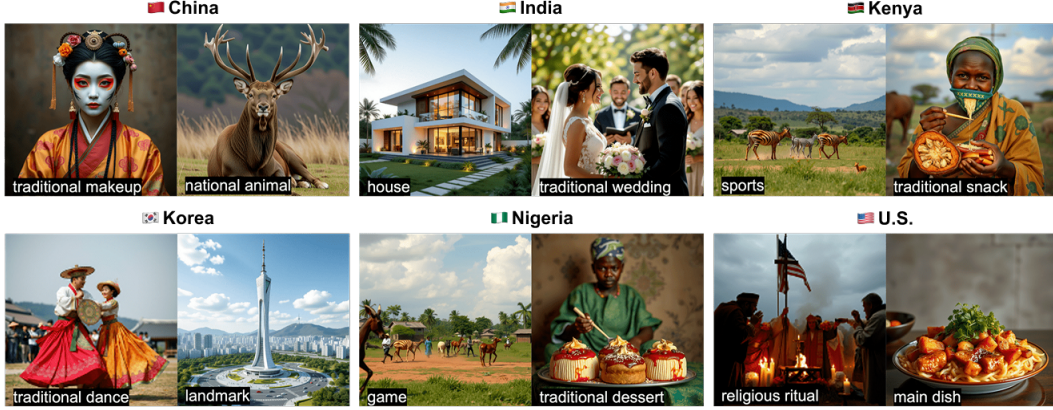


Figure 1: Representative cultural biases in T2I generations across six countries. Examples include Chinese–Japanese aesthetic conflation, mis-styled Indian weddings, Kenyan wildlife stereotypes, Korean attire misidentification, Nigerian safari mislocalization, and U.S. cultural miscues in food and religious ritual. Images are from FLUX.1 [schnell] fp8 and HiDream-I1-Dev.

1 Introduction

Text-to-image (T2I) and image-to-image (I2I) generative models have advanced rapidly in photorealism and controllability [1, 2, 3]. Yet T2I models exhibit systematic cultural bias rooted in imbalanced, skewed training data, where some regions and communities are over-represented while others are under-represented [4, 5]. This yields stereotyped or inaccurate portrayals of underrepresented groups, as shown in Figure 1. I2I editing is often used as a practical remedy by inserting or adjusting culture-specific elements [6], but errors persist across regions and frequently resurface over multiple edits [7, 8]. As a result, the burden of cultural correction shifts to users. We therefore focus on the *I2I editing loop*: a sequence of edits intended to align an image with a target culture, used to examine whether bias persists, attenuates, or reappears across iterations.

Turning to evaluation, cultural fidelity remains challenging to measure. Distributional or general-alignment metrics such as Fréchet Inception Distance (FID) [9] and CLIPScore [10] do not capture culture-specific attributes. Emerging cultural benchmarks [8, 11] help, but often rely heavily on human studies or use only author-generated synthetic images, limiting generalization [12].

We analyze five representative model families (Table 1) by first constructing a T2I base image set spanning six countries: China, India, Kenya, Korea, Nigeria, and the United States (U.S.). The set is balanced across eight cultural categories (Architecture, Art, Events, Fashion, Food, Landscape, People, Wildlife), further stratified into 36 subcategories (Table 2). Our prompts are *era-aware*, comprising *traditional*, *modern*, and *era-agnostic* variants, which allows us to probe temporal sensitivity in cultural understanding. To the best of our knowledge, this is the first study to systematically evaluate era-aware cultural competence in image models. Building on the base image set, we evaluate I2I cultural adaptation using three complementary experiments: (1) Multi-Loop Edit: five sequential edits that test whether the model preserves context while progressively improving cultural consistency; (2) Attribute Addition: stepwise insertion of five distinct culture-specific attributes to quantify cumulative competence and interference; and (3) Cross-Country Restylization: coherent transfer from a source country to a target country to assess adaptability and generalization.

For evaluation, we adopt automatic metrics—CLIPScore [10], DreamSim [13], and Aesthetic Score [14]—complemented by culture-centered assessments. We compare automatic metrics against two complementary evaluations: a VQA-based culture-aware metric and a human evaluation. This approach allows us to analyze the agreement, gaps, and limitations between automatic evaluations and human perception.

Our controlled prompting method is designed to reduce variability from user phrasing and isolate model-side cultural tendencies, framing this study as a diagnostic stress test rather than a full model of all real-world usage. We also conduct an additional bias analysis focusing on gender and skin-

tone stereotyping that persists across generative models (see the extended version of this paper¹). These findings point to a concrete need: curated and balanced training data, explicit training-time debiasing/regularization, and standardized fairness benchmarks to enable equitable and responsible deployment.

Our contributions are as follows:

1. **Experimental design.** A geographically balanced, multi-domain schema (6 countries; 8 categories, 36 subcategories) and 3 complementary protocols (Multi-Loop Edit, Attribute Addition, Cross-Country Restylization) that contrast inherent T2I bias with I2I editing capability, including an era-aware prompt design to assess temporal awareness (an underexplored evaluation setting).
2. **Dataset release.** Public release of the complete image set, together with prompts and model/execution configurations, enabling downstream cultural analyses without re-running generation/editing or model reconfiguration.
3. **Evaluation framework.** An open-source evaluation platform that pairs automated metrics with culture-aware assessments and human studies, enabling triangulation across automated, culture-aware, and human judgments; we benchmark and validate metric behavior against human results to surface agreements and discrepancies.

2 Related Work

2.1 Generative Image Models

Generative image models rely on two closely related paradigms: T2I generation and I2I editing. Latent diffusion models [15, 2] set the T2I standard for fidelity and controllability, yet often miss fine-grained cultural fidelity. In response, I2I frameworks increase edit precision via context-aware conditioning, multi-step guidance, and cross-attention manipulation (e.g., FLUX.1 Kontext [16], HiDream-I1 [17], Qwen-Image-Edit [18], NextStep-1 [19]). These methods preserve locality and enable targeted attributes but are typically optimized for realism and generic prompt compliance rather than culture-specific correctness. Prior studies [20, 21, 22] show that stronger editing control alone does not mitigate entrenched biases: models can encode or even amplify stereotypes when data and objectives are not culture-aware. We therefore compare T2I baselines and their I2I editing loops to quantify initial cultural bias and the bias-correction burden placed on users.

2.2 Text-Image Datasets

The foundation of modern generative image models rests on large-scale text-image datasets. Early efforts like MS-COCO [23] and ImageNet [24] were predominantly Western-centric. Subsequent expansions, including LAION-5B [25], achieved unprecedented scale but inherently inherited severe cultural and geographical biases, notably English dominance and uneven global representation. While datasets like Dollar Street [26] and WIT [27] attempted to target diversity, their scope was limited by inherent editorial constraints. More recent benchmarks, including CCUB [28] and SCoFT [29], have focused on directly assessing cultural fidelity. However, these resources primarily prioritize concept coverage over the representation of nuanced cultural practices, ceremonies, or daily life contexts [30]. This persistent gap—where existing datasets remain ill-suited for assessing authentic cultural representation—motivates the creation of our targeted evaluation framework, which is the direct focus of this work.

2.3 Evaluation Metrics in Cultural Image Generation

The evolution of generative image models mandates corresponding advancements in output evaluation. Traditional metrics mainly assess image quality (e.g., FID [9], DreamSim [13], Aesthetic Score [14]) and text-image alignment (e.g., CLIPScore [10] and DinoScore [31]). Recent improvements leverage human preference-trained models [32] and large language model (LLM)-based metrics [33]. However, these approaches primarily capture realism and faithfulness to prompts, while neglecting cultural and contextual fidelity [8, 34]. Recent benchmarks such as CUBE [8], CULTDIFF [11], and CulturalFrames [30] aim to evaluate cultural alignment and bias, yet often depend on human judgments or limited internal datasets, constraining generalizability. To overcome this gap, we propose

¹An extended version with all technical appendices is available at <https://arxiv.org/abs/2510.20042>.

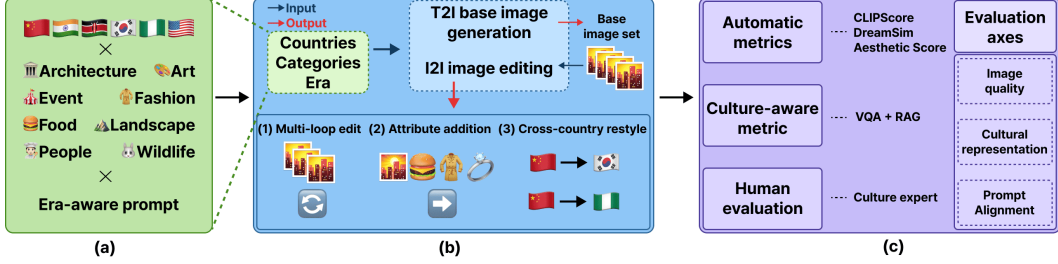


Figure 2: Overall framework overview. (a) Schema inputs: six countries, eight categories, and three era-aware prompts. (b) Experimental pipeline: T2I base generation and three I2I editing studies. (c) Multi-layered evaluation: integrating automatic, culture-aware metrics, and human evaluation.

a comprehensive framework for evaluating authentic cultural representation in general generative outputs.

3 Experiment Design for Cultural Bias Evaluation

3.1 Experimental Setup

We design two complementary studies: a T2I study for bias via image generation, and an I2I study for competence via editing culturally salient details. We evaluate across six geographically diverse countries (China, India, Kenya, Korea, Nigeria, U.S.), selected to span multiple regions while keeping expert evaluation feasible, addressing underrepresentation in training data [30]. The overall framework is illustrated in Figure 2.

Category schema. Building on prior works [8, 35], our schema defines eight top-level categories (36 subcategories) covering both tangible artifacts (e.g., architecture, cuisine) and intangible practices (e.g., rituals, festivals), roles, and settings. This granular system probes cultural competence by enabling era-aware prompting and controlled cross-country comparisons to evaluate breadth and depth, rather than aesthetics alone systematically. The full schema is shown in Table 2.

Models. We conducted T2I and I2I experiments using state-of-the-art open-source models. Models were selected to capture capability diversity and ensure reproducibility. To our knowledge, none of these models is officially reported by their developers as targeting any specific region. The full list of models is summarized in Table 1 with all settings detailed in the extended version.

Evaluation metrics. We evaluate models using two complementary dimensions: general-purpose and culture-aware. For general-purpose, we adopt CLIPScore, DreamSim, and Aesthetic Score to measure semantic alignment, cumulative edit distance, and visual quality. For culture-aware evaluation, we extend VQA- and retrieval-augmented generation (RAG)-based approaches [35, 36] by retrieving Wikipedia context via FAISS [37] and generating yes/no questions and answer with Qwen2.5-0.5B and Qwen2-VL-7B [38, 39]. Human ratings serve as the primary reference, enabling analysis of consistency and divergence between automatic and culture-aware metrics.

Prompt design. To disentangle temporal stylization from cultural identity, we use three prompt modes for each country/category/subcategory:

Table 1: Models used in T2I and I2I experiments.

Family	T2I	I2I
Stable Diffusion [2]	Stable Diffusion 3.5 Medium	Stable Diffusion 3.5 Medium
FLUX.1 [16]	FLUX.1 [schnell] fp8	FLUX.1 Kontext [dev]
HiDream [17]	HiDream-I1-Dev	HiDream-E1.1
Qwen-Image [18]	Qwen-Image	Qwen-Image-Edit
NextStep [19]	NextStep-1-Large	NextStep-1-Large-Edit

- Traditional: “Traditional {subcategory} in {Country}, photorealistic.”
- Modern: “Modern {subcategory} in {Country}, photorealistic.”
- Era-agnostic: “{Subcategory} in {Country}, photorealistic.”

This era-aware prompting allows us to examine how cultural depictions vary across specified and unspecified eras, revealing era-sensitive biases.

3.2 Text-to-Image (T2I) Experiment

We construct a standardized base image set to expose model-internal cultural priors under controlled text-only prompting. For each country and subcategory, we issue era-aware prompts (traditional, modern, and era-agnostic) while fixing all sampling parameters within each model family (Table 1, T2I column). This yields a comparable collection of generations across all settings, which is subsequently used for distributional analyses, traditional-modern scoring, and seeding the I2I editing experiments in Section 3.3. Fixing prompts and sampling settings isolates differences attributable to the models rather than to prompt phrasing or parameter drift.

3.3 Image-to-Image (I2I) Experiment

We evaluate cultural editing competence with three experimental designs using open-source I2I models—(Table 1, I2I column)—reflecting contemporary editing paradigms while remaining reproducible.

3.3.1 Multi-loop edit: cultural consistency under iterative edits

To minimize cross-model domain shifts, base images are drawn from a single T2I model family. We then apply the instruction “Change the image to represent {era} {subcategory} in {Country}” for five successive rounds. This iterative experiment tests stability and path dependence in cultural editing—whether iterations converge toward culturally faithful attributes or drift by amplifying stereotypes.

3.3.2 Attribute addition: compositional cultural grounding

To isolate composition effects, we begin with a neutral canvas: a genderless green mannequin on a white background devoid of cultural cues. For each country/model, a fixed five-step sequence controls (1) background, (2) local-script text rendering, (3) food, (4) clothing, and (5) traditional accessories. This sequence allows us to assess, within a single experiment, the model’s understanding across multiple facets commonly flagged as challenging in cultural benchmarks. Full stepwise prompts are provided in the extended version.

3.3.3 Cross-country restyle: style transfer across countries

We use T2I base images from HiDream-I1-Dev (selected for strong T2I fidelity) and test whether I2I models can restyle an entire scene into a target country’s aesthetic while preserving the subject’s identity and pose. We employ a single minimal prompt—“Transform this image into the {CountryAdj} style.”—and report results under this condition; in preliminary trials, adding

Table 2: Category schema used in our experiments.

category	subcategories	category	subcategories
Architecture	House, Landmark	Landscape	City, Countryside, Nature
Art	Dance, Painting, Sculpture	Fashion	Accessories, Clothing, Makeup
Food	Beverage, Dessert, Main dish, Snack, Staple food	Wildlife	Animal, Plant
People	Daily life, Athlete, Bride and groom, Celebrity, Chef, Doctor, Farmer, Model, President, Soldier, Student, Teacher		



Figure 3: Comparative samples for U.S. (top) and country-agnostic (bottom) prompts across two models. Left image: FLUX.1 [schnell] fp8; right image: HiDream-I1-Dev. Within each model, panels show (from left) *Bride and groom*, *Chef*, and *Farmer*. The close correspondence between rows illustrates the US-like default of country-agnostic prompts.

era/context/attribute cues did not materially improve cultural plausibility, so we do not vary prompt specificity further. This setup supports (i) qualitative audits of cultural appropriateness and context/era consistency and (ii) inspection of edit stability across multi-loop steps.

3.4 Expert Human Evaluation

Our human evaluation is intentionally minimal to highlight a core research question: can existing automatic metrics substitute for human evaluation on cultural content? Specifically, we compare CLIPScore against human prompt-alignment ratings and Aesthetic Score against human-rated image quality. We complement automatic metrics with an expert-group human evaluation, run on our web platform (ECB Human Survey) under a unified protocol. Raters evaluate only their own country (emic expertise), ensuring cultural authenticity in judgments. For each prompt, four candidates (step base/1/3/5) are displayed side-by-side. Raters assign two scores: 1-5 Likert scores for *Image Quality* and *Cultural Representation*. These two components are averaged to form the Human Quality Score (HQS). Raters also select *Best/Worst*, following recent practices in cultural assessment and editing evaluation [12, 40]. Image sets span eight categories across five models (Table 1). Each rater completes six tasks: five I2I-loop evaluations and one attribute addition evaluation. Expert raters require insider knowledge and language proficiency. Operational details, including the distribution of our expert raters, are provided in the extended version.

4 Results

4.1 Text-to-Image (T2I)

We use a unified embedding backbone across analyses. Let $x_i \in \mathbb{R}^d$ denote the CLIP image embedding of sample i generated by one of $M=5$ T2I models, and let $c(i) \in \mathcal{C}$ be its country tag (*country-agnostic* means no explicit tag). When clustering is required, we perform model-wise PCA to two components and run k -means with K_m clusters. All mathematical definitions (cluster-proportion vectors, Jensen-Shannon divergence (JSD), the proximity metrics, and the traditional-modern leaning score with its summaries) are given in the extended version of this paper.

4.1.1 Distributional Proximity and Cultural Defaults

For each model m and country c , we compute cluster proportions and measure pairwise country similarity using a JSD-based proximity score, which is then averaged across models (full definitions in the extended version). Across all five models, the strongest proximity is U.S. $\leftrightarrow$ country-agnostic ($\bar{h}=0.892$, 95% CI [0.844, 0.931]), followed by China $\leftrightarrow$ Korea (0.888, [0.835, 0.932]) and Kenya $\leftrightarrow$ Nigeria (0.864, [0.811, 0.911]). A meta-analysis indicates negligible between-model heterogeneity ($\tau^2 \approx 0$ except Kenya–Nigeria $\tau^2 \approx 8.5 \times 10^{-5}$), indicating stable effects. These findings support a robust U.S.-like default for country-agnostic prompts and regional clustering in East Asia and Sub-Saharan Africa. Model-level nearest-neighbor votes confirm China–Korea and Kenya–Nigeria as

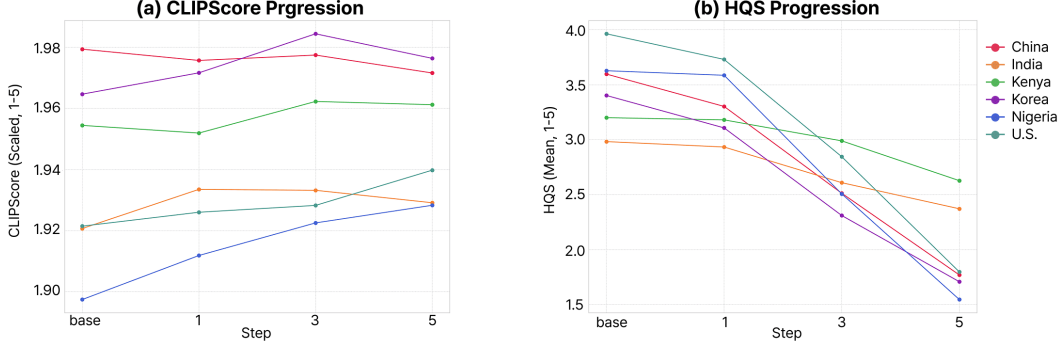


Figure 4: Divergence between Automated and Human Judgment in Iterative Editing. (a) CLIPScore trajectories remain largely stable or modestly increase from the base step to step 5. (b) Human Quality Score (HQS) sharply declines across all countries. This pronounced divergence highlights the failure of traditional automatic metrics to track the perceptible cultural degradation that human raters consistently penalize.

mutual neighbors, while India often aligns with U.S./country-agnostic, suggesting partial assimilation. See Figure 3 for a visual comparison: U.S. and country-agnostic prompts produce almost identical samples, qualitatively presenting the U.S.-like default.

4.1.2 Traditional–Modern Leaning under Country-agnostic Prompts

Image-level traditional-modern scores are computed per model and aggregated to country means (definitions in the extended version), and their significance is evaluated using within-category permutation with BH-FDR across countries. The cross-country dispersion of $\{\bar{s}_c\}$ is significant in every model (four models: $p=0.001$; Qwen-Image: $p=0.002$), indicating systematic differences under country-agnostic prompts. The dominant pattern is a consistent *modern* lean for the U.S. and country-agnostic across all five models ($\bar{s}_c < 0$, per-model $p \leq 0.006$, FDR-significant). Other effects are small and model-dependent: Kenya is traditional-leaning only in FLUX.1 [schnell] fp8 ($p=0.001$, FDR-significant), while India, Nigeria, Korea, and China are mixed or non-significant (typically $q > 0.1$). Category composition largely explains the country means: *clothing* and *sport* consistently yield modern (negative) scores across models and countries—especially for the U.S. and country-agnostic—whereas *house*, *religious ritual*, and *landmark* generally yield traditional (positive) scores for most countries but often flip to modern for the U.S. and country-agnostic. These patterns indicate that the modern default arises from systematic topic/style composition rather than sampling noise. Per-country p -values and BH-FDR-adjusted q -values are reported in the extended version.

4.2 Image-to-Image (I2I)

4.2.1 Multi-Loop Edit

Following the multi-loop experiment in Section 3.3.1, we report country-level averages across models for five successive I2I edits. Our findings reveal a substantial gap between automatic metrics and human judgment. As shown in Figure 4, CLIPScore remains largely flat or slightly increases across steps, whereas human-related HQS sharply declines for most model-country pairs. This divergence indicates that with continued edits, culturally salient cues—such as context and era—gradually erode, yielding images that appear more aligned with prompts yet remain culturally distorted.

Other automatic metrics, including Aesthetic Score and DreamSim, exhibit modest or decreasing trends but fail to capture culture-specific degradation. In contrast, our culture-aware metric shows a stepwise decline consistent with HQS, demonstrating its sensitivity to cultural loss. Figure 5 further illustrates this alignment: the agreement with human selection reaches 73.8% for best images and 83.7% for worst images.

Overall, while traditional automatic metrics remain stable or even improve despite perceptible cultural deterioration, our culture-aware metric aligns more closely with human perception. These results suggest that culturally sensitive editing becomes unreliable under iterative I2I, where standard metrics

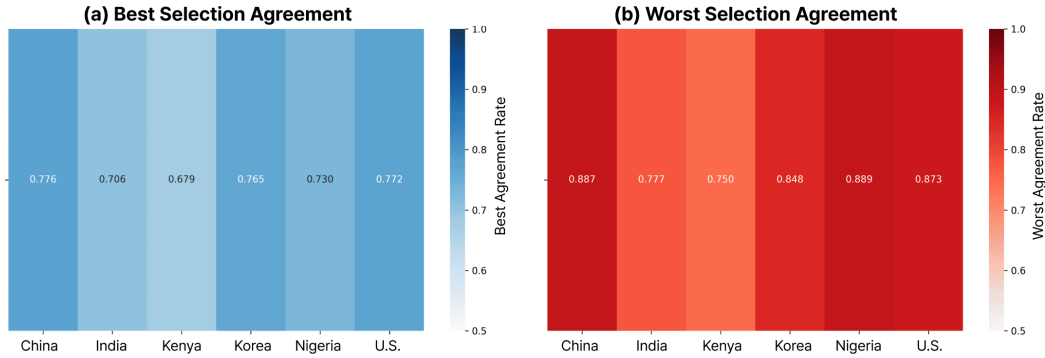


Figure 5: Alignment of the Culture-aware Metric with Human Judgment (All Models Averaged). (a) The agreement rate for *Best Selection* is high across all countries, averaging 73.8%. (b) The agreement rate for *Worst Selection* is consistently higher, averaging 83.7%. This high alignment demonstrates that our extended culture-aware metric successfully tracks human preference for unedited states and penalizes pronounced cultural erosion. Detailed stepwise score changes are provided in the extended version.

may obscure degradation that both human evaluators and our extended metric reveal. Detailed per-model and per-country analyses, together with additional qualitative examples, are provided in the extended version.

4.2.2 Attribute Addition

Following the attribute-addition protocol in Section 3.3.2, we test whether models can compose culture-specific elements step by step. We sequentially add five attributes—background, local-script text, food, clothing, and accessories—and visualize representative progressions in Figure 6. For evaluation, the expert group in Section 3.4 rated each step on three dimensions (Likert): image quality, prompt alignment, and cultural fidelity. The results are as follows: The sequential attribute-addition task quantified compositional failure, revealing a consistent decline in image quality (IQ) due to the compounding of edits. Crucially, while FLUX.1 Kontext [dev] demonstrated superior overall IQ, Qwen-Image-Edit achieved higher prompt alignment in complex compositional tasks (Food/Clothing), suggesting a fundamental tradeoff between strict object control and visual quality. This analysis highlights two critical failure points: the Text attribute (sharp alignment declines from gibberish) and the final Accessory attribute (lowest IQ scores), indicating persistent difficulty rendering complex, culturally specific details.

4.2.3 Cross-Country Restylization

Using the single minimal prompt (Section 3.3.3), we qualitatively assess two I2I models; Figure 7 shows the outcomes. Both models preserve layout, composition, and pose across multi-loop steps, but culture-relevant details drift: localized attributes (attire, objects, vernacular architecture, signage) remain under-edited or degrade. Qwen-Image-Edit often substitutes symbolic or palette shifts associated with the target country for true localized changes, yielding surface “signals” rather than context- or era-consistent edits. When targeting non-U.S. countries (particularly those from Africa and Asia), subjects frequently retain their original phenotype (race/skin tone), indicating weak identity adaptation. We also observe a style asymmetry: non-U.S. targets are commonly rendered in drawing/painting styles, whereas U.S. targets remain photorealistic. Overall, under minimal prompting, current editing models favor superficial markers and structural conservation over culture- and era-faithful transformations. Additional qualitative examples are included in the extended version.

5 Discussion and Limitations

Discussion. Our experiments show substantial cultural bias in current image generators, with three dominant patterns. (i) U.S.-like default: country-agnostic prompts produce U.S.-like, modern-leaning outputs (Figure 3); stable regional pairings (e.g., China-Korea; Kenya-Nigeria) indicate that category

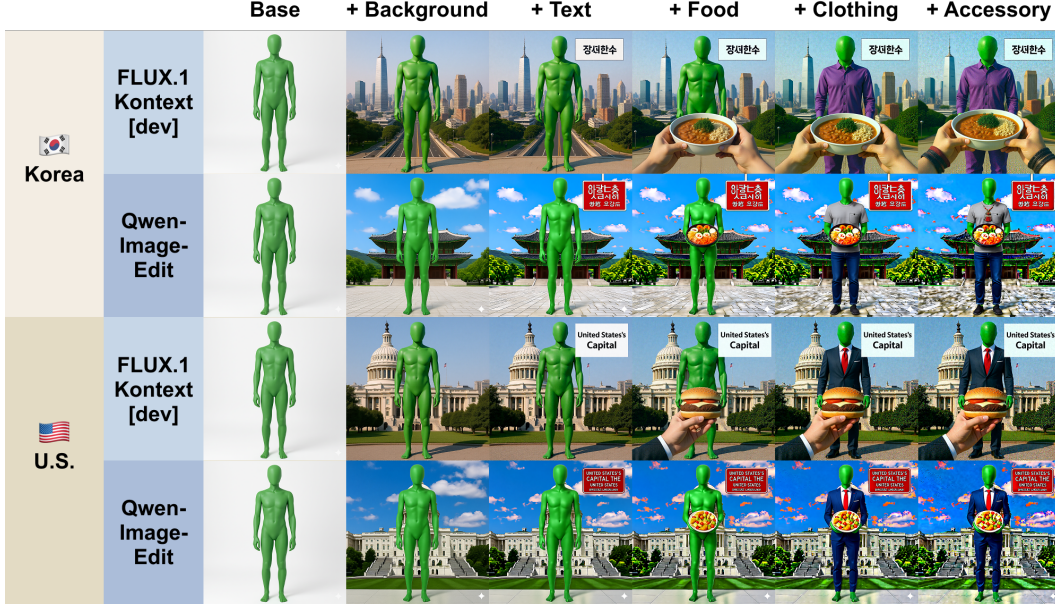


Figure 6: Representative stepwise attribute additions for Korea and the U.S. Columns progress from the base image to background, text, food, clothing, and accessory, each added cumulatively. Rows denote the model; the top block uses Korea prompts and the bottom block U.S. prompts.

mix and style priors—not noise—drive cross-country gaps. (ii) Metric-human gap in iterative I2I: conventional metrics (e.g., CLIPScore) stay flat or rise while human-rated cultural quality drops (Figure 4); our culture-sensitive metric tracks human Best/Worst choices (Figure 5), revealing metric drift. (iii) Shortcut editing: models “signal” culture via superficial cues (flag overlays, palette swaps), preserve source identity when targeting non-U.S. countries, and render non-U.S. scenes in painterly styles while keeping U.S. scenes photorealistic (Figure 6, 7). These cultural effects co-occur with occupation-level demographic skews—male dominance in several roles, female skew in caregiving/aesthetic roles, and light-tone prevalence under neutral prompts (see the extended version).

Limitations. Our I2I experiments start from within-family T2I baselines to isolate intra-model behavior and reduce cross-family confounds, which limits claims about transfer across model families. We use generic prompts (no per-model engineering), so results reflect out-of-the-box behavior rather than tuned best case. Scale (five models, six countries) is bounded by human-evaluation cost and compute constraints—multi-loop editing and large-batch inference require substantial GPU resources. A further limitation is the unit of analysis: we use country-level labels as a proxy for “culture” (sovereign states, e.g., U.S., China), not subnational units. Prior work cautions that aligning culture with geopolitical borders obscures within-country heterogeneity, minority communities, and transnational/diaspora groups [12, 11]; this is salient for the U.S. (immigration-shaped, strong regional variation) and for China (many officially recognized minority groups). We also acknowledge the risk of cultural essentialism; to mitigate this, we use era-aware prompts, attribute-addition tests, and flexible contextual cues rather than prescriptive symbols. Nevertheless, intra-country diversity, the potential for stereotyping, and concerns about cultural offensiveness warrant further study. Finally, automated components used in evaluation (e.g., VQA/LLM-assisted steps) can introduce their own biases, so results should be interpreted with care [11]. Future work should adopt finer-grained and potentially multi-label groupings (subnational regions, language/community groups) and audit automated tools, as recommended by prior studies [12, 11].

6 Conclusion and Future Work

We presented a structured evaluation of cultural bias across T2I and I2I using six countries, an 8-category/36-subcategory schema, era-aware prompts (*traditional/modern/era-agnostic*), and three



Figure 7: Representative cross-country restyling results for China and Nigeria; columns are ordered by geographic proximity to the base country—closest on the left, farthest on the right. Rows indicate the model; the top block uses a China base image and the bottom block a Nigeria base image.

protocols (Multi-Loop Edit, Attribute Addition, Cross-Country Restylization). Our open platform pairs standard automatic metrics with a culture-sensitive metric and a web-based human study workflow, enabling triangulation and model-level diagnostics. Empirically, country-agnostic prompting defaults to U.S.-like, modern-leaning outputs; iterative I2I can erode cultural fidelity while conventional metrics obscure the decline; and editing pipelines often rely on superficial cues rather than culture-consistent, context-preserving changes. Progress hinges on two practical directions. *Finer granularity*: report beyond country tags—include subnational (state/province/city) and community/language groups—and use simple stratified reporting so diverse communities are not collapsed into a single label. *Stronger training/data signals*: curate balanced datasets; add era-aware, context-preserving objectives to both generation and editing; and explicitly penalize shortcut cues that fake culture (e.g., flag overlays, generic palettes, style flips that ignore people/objects/setting). We release all images, prompts, and configurations to support reproducible follow-up studies along these directions.

Acknowledgments and Disclosure of Funding

We would like to express our sincere gratitude to the following contributors for their valuable support and collaboration throughout this work: Parth Maheshwari, Abdullahi Abdulrahman, Ge Fang, Dusillah Dullo, Alice Etori, Abdullahi Adavize Ismaila, SoHee Yoon, Jamin Lee, Ikbum Park, Taeseo Kim, and Hyowon Choi. This work is in part supported by NSF IIS-2112633 and J.P. Morgan. Minki Hong, Sieun Choi, and Jihie Kim are supported by the MSIT (Ministry of Science and ICT) of Korea under the Global Research Support Program in the Digital Field (RS-2024-00426860) and the Artificial Intelligence Convergence Innovation Human Resources Development (IITP-2026-RS-2023-00254592), supervised by the IITP (Institute for Information & Communications Technology Planning & Evaluation).

References

- [1] Chitwan Saharia, William Chan, Saurabh Saxena, Lala Li, Jay Whang, Emily L Denton, Kamyar Ghasemipour, Raphael Gontijo Lopes, Burcu Karagol Ayan, Tim Salimans, et al. Photorealistic text-to-image diffusion models with deep language understanding. *Advances in neural information processing systems*, 35:36479–36494, 2022.
- [2] Patrick Esser, Sumith Kulal, Andreas Blattmann, Rahim Entezari, Jonas Müller, Harry Saini, Yam Levi, Dominik Lorenz, Axel Sauer, Frederic Boesel, et al. Scaling rectified flow transformers for high-resolution image synthesis. In *Forty-first international conference on machine learning*, 2024.
- [3] Zineb Sordo, Eric Chagnon, and Daniela Ushizima. A review on generative ai for text-to-image and image-to-image generation and implications to scientific images. *arXiv preprint arXiv:2502.21151*, 2025.
- [4] Siddharth Kandwal and Vibha Nehra. A survey of text-to-image diffusion models in generative ai. In *2024 14th International Conference on Cloud Computing, Data Science & Engineering (Confluence)*, pages 73–78. IEEE, 2024.
- [5] Yixin Wan, Arjun Subramonian, Anaelia Ovalle, Zongyu Lin, Ashima Suvarna, Christina Chance, Hritik Bansal, Rebecca Pattichis, and Kai-Wei Chang. Survey of bias in text-to-image generation: Definition, evaluation, and mitigation. *arXiv preprint arXiv:2404.01030*, 2024.
- [6] Simran Khanuja, Sathyanarayanan Ramamoorthy, Yueqi Song, and Graham Neubig. An image speaks a thousand words, but can everyone listen? on image transcreation for cultural relevance. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pages 10258–10279, 2024.
- [7] Bingshuai Liu, Longyue Wang, Chenyang Lyu, Yong Zhang, Jinsong Su, Shuming Shi, and Zhaopeng Tu. On the cultural gap in text-to-image generation. *arXiv preprint arXiv:2307.02971*, 2023.
- [8] Nithish Kannen, Arif Ahmad, Marco Andreetto, Vinodkumar Prabhakaran, Utsav Prabhu, Adji Bousso Dieng, Pushpak Bhattacharyya, and Shachi Dave. Beyond aesthetics: cultural competence in text-to-image models. In *Proceedings of the 38th International Conference on Neural Information Processing Systems*, pages 13716–13747, 2024.
- [9] Martin Heusel, Hubert Ramsauer, Thomas Unterthiner, Bernhard Nessler, and Sepp Hochreiter. Gans trained by a two time-scale update rule converge to a local nash equilibrium. *Advances in neural information processing systems*, 30, 2017.
- [10] Jack Hessel, Ari Holtzman, Maxwell Forbes, Ronan Le Bras, and Yejin Choi. Clipscore: A reference-free evaluation metric for image captioning. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 7514–7528, 2021.
- [11] Zahra Bayramli, Ayhan Suleymanzade, Na Min An, Huzama Ahmad, Eunsu Kim, Junyeong Park, James Thorne, and Alice Oh. Diffusion models through a global lens: Are they culturally inclusive? *arXiv preprint arXiv:2502.08914*, 2025.
- [12] Shravan Nayak, Mehar Bhatia, Xiaofeng Zhang, Verena Rieser, Lisa Anne Hendricks, Sjoerd van Steenkiste, Yash Goyal, Aishwarya Agrawal, et al. Culturalframes: Assessing cultural expectation alignment in text-to-image models and evaluation metrics. *arXiv preprint arXiv:2506.08835*, 2025.
- [13] Stephanie Fu, Netanel Y Tamir, Shobhita Sundaram, Lucy Chai, Richard Zhang, Tali Dekel, and Phillip Isola. Dreamsim: learning new dimensions of human visual similarity using synthetic data. In *Proceedings of the 37th International Conference on Neural Information Processing Systems*, pages 50742–50768, 2023.
- [14] Heng Huang, Xin Jin, Yaqi Liu, Hao Lou, Chaoen Xiao, Shuai Cui, Xining Li, and Dongqing Zou. Predicting scores of various aesthetic attribute sets by learning from overall score labels. In *Proceedings of the 2nd International Workshop on Multimedia Content Generation and Evaluation: New Methods and Practice*, pages 63–71, 2024.
- [15] Robin Rombach, Andreas Blattmann, Dominik Lorenz, Patrick Esser, and Björn Ommer. High-resolution image synthesis with latent diffusion models. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 10684–10695, 2022.
- [16] Black Forest Labs, Stephen Batifol, Andreas Blattmann, Frederic Boesel, Saksham Consul, Cyril Diagne, Tim Dockhorn, Jack English, Zion English, Patrick Esser, et al. Flux. 1 konteks: Flow matching for in-context image generation and editing in latent space. *arXiv preprint arXiv:2506.15742*, 2025.

- [17] Qi Cai, Jingwen Chen, Yang Chen, Yehao Li, Fuchen Long, Yingwei Pan, Zhaofan Qiu, Yiheng Zhang, Fengbin Gao, Peihan Xu, et al. Hidream-1l: A high-efficient image generative foundation model with sparse diffusion transformer. *arXiv preprint arXiv:2505.22705*, 2025.
- [18] Chenfei Wu, Jiahao Li, Jingren Zhou, Junyang Lin, Kaiyuan Gao, Kun Yan, Sheng-ming Yin, Shuai Bai, Xiao Xu, Yilei Chen, et al. Qwen-image technical report. *arXiv preprint arXiv:2508.02324*, 2025.
- [19] NextStep Team, Chunrui Han, Guopeng Li, Jingwei Wu, Quan Sun, Yan Cai, Yuang Peng, Zheng Ge, Deyu Zhou, Haomiao Tang, et al. Nextstep-1: Toward autoregressive image generation with continuous tokens at scale. *arXiv preprint arXiv:2508.10711*, 2025.
- [20] Federico Bianchi, Pratyusha Kalluri, Esin Durmus, Faisal Ladhak, Myra Cheng, Debora Nozza, Tatsunori Hashimoto, Dan Jurafsky, James Zou, and Aylin Caliskan. Easily accessible text-to-image generation amplifies demographic stereotypes at large scale. In *Proceedings of the 2023 ACM conference on fairness, accountability, and transparency*, pages 1493–1504, 2023.
- [21] Ranjita Naik and Besmira Nushi. Social biases through the text-to-image generation lens. In *Proceedings of the 2023 AAAI/ACM Conference on AI, Ethics, and Society*, pages 786–808, 2023.
- [22] Aditya Chinchure, Pushkar Shukla, Gaurav Bhatt, Kiri Salij, Kartik Hosanagar, Leonid Sigal, and Matthew Turk. Tibet: Identifying and evaluating biases in text-to-image generative models. In *European Conference on Computer Vision*, pages 429–446. Springer, 2024.
- [23] Tsung-Yi Lin, Michael Maire, Serge Belongie, James Hays, Pietro Perona, Deva Ramanan, Piotr Dollár, and C Lawrence Zitnick. Microsoft coco: Common objects in context. In *European conference on computer vision*, pages 740–755. Springer, 2014.
- [24] Jia Deng, Wei Dong, Richard Socher, Li-Jia Li, Kai Li, and Li Fei-Fei. Imagenet: A large-scale hierarchical image database. In *2009 IEEE conference on computer vision and pattern recognition*, pages 248–255. Ieee, 2009.
- [25] Christoph Schuhmann, Romain Beaumont, Richard Vencu, Cade Gordon, Ross Wightman, Mehdi Cherti, Theo Coombes, Aarush Katta, Clayton Mullis, Mitchell Wortsman, et al. Laion-5b: An open large-scale dataset for training next generation image-text models. *Advances in neural information processing systems*, 35:25278–25294, 2022.
- [26] William Gaviria Rojas, Sudnya Diamos, Keertan Kini, David Kanter, Vijay Janapa Reddi, and Cody Coleman. The dollar street dataset: Images representing the geographic and socioeconomic diversity of the world. *Advances in Neural Information Processing Systems*, 35:12979–12990, 2022.
- [27] Krishna Srinivasan, Karthik Raman, Jiecao Chen, Michael Bendersky, and Marc Najork. Wit: Wikipedia-based image text dataset for multimodal multilingual machine learning. In *Proceedings of the 44th international ACM SIGIR conference on research and development in information retrieval*, pages 2443–2449, 2021.
- [28] Zhixuan Liu, Youeun Shin, Beverley-Claire Okogwu, Youngsik Yun, Lia Coleman, Peter Schaldenbrand, Jihie Kim, and Jean Oh. Towards equitable representation in text-to-image synthesis models with the cross-cultural understanding benchmark (ccub) dataset. *arXiv preprint arXiv:2301.12073*, 2023.
- [29] Zhixuan Liu, Peter Schaldenbrand, Beverley-Claire Okogwu, Wenxuan Peng, Youngsik Yun, Andrew Hundt, Jihie Kim, and Jean Oh. Scoft: Self-contrastive fine-tuning for equitable image generation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 10822–10832, 2024.
- [30] Shravan Nayak, Kanishk Jain, Rabiul Awal, Siva Reddy, Sjoerd Steenkiste, Lisa Hendricks, Karolina Stanczak, and Aishwarya Agrawal. Benchmarking vision language models for cultural understanding. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pages 5769–5790, 2024.
- [31] Nataniel Ruiz, Yuanzhen Li, Varun Jampani, Yael Pritch, Michael Rubinstein, and Kfir Aberman. Dream-booth: Fine tuning text-to-image diffusion models for subject-driven generation. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 22500–22510, 2023.
- [32] Xiaoshi Wu, Yiming Hao, Keqiang Sun, Yixiong Chen, Feng Zhu, Rui Zhao, and Hongsheng Li. Human preference score v2: A solid benchmark for evaluating human preferences of text-to-image synthesis. *CoRR*, 2023.
- [33] Zhiqiu Lin, Deepak Pathak, Baiqi Li, Jiayao Li, Xide Xia, Graham Neubig, Pengchuan Zhang, and Deva Ramanan. Evaluating text-to-visual generation with image-to-text generation. *CoRR*, 2024.

- [34] Muna Numan Said, Aarib Zaidi, Rabia Usman, Sonia Okon, Praneeth Medepalli, Kevin Zhu, Vasu Sharma, and Sean O’Brien. Deconstructing bias: A multifaceted framework for diagnosing cultural and compositional inequities in text-to-image generative models. *arXiv preprint arXiv:2505.01430*, 2025.
- [35] 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. In *Proceedings of the 38th International Conference on Neural Information Processing Systems*, pages 11479–11505, 2024.
- [36] Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, et al. Retrieval-augmented generation for knowledge-intensive nlp tasks. *Advances in neural information processing systems*, 33:9459–9474, 2020.
- [37] Jiaang Li, Yifei Yuan, Wenyan Li, Mohammad Aliannejadi, Daniel Hershcovich, Anders Søgaard, Ivan Vulić, Wenxuan Zhang, Paul Pu Liang, Yang Deng, et al. Ravenea: A benchmark for multimodal retrieval-augmented visual culture understanding. *arXiv preprint arXiv:2505.14462*, 2025.
- [38] Qwen, :, 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.
- [39] Qwen Team et al. Qwen2 technical report. *arXiv preprint arXiv:2407.10671*, 2:3, 2024.
- [40] Zhuoying Li, Zhu Xu, Yuxin Peng, and Yang Liu. Balancing preservation and modification: A region and semantic aware metric for instruction-based image editing. *arXiv preprint arXiv:2506.13827*, 2025.