

Simran Kaur

Beyond the Western Gaze:

A Critical Analysis of MidJourney's West-Centric Lens on South Asian Representation

Keywords: AI Bias, South Asian Representation, Prompt Engineering, Indigenous Epistemologies, Decolonial Theory

Abstract

This paper examines systematic biases in AI image generation, specifically focusing on MidJourney's representation of South Asian subjects through a decolonial theoretical framework. Through analysis of over 1,000 AI-generated images and evaluation by a panel of twelve experts including media scholars, photographers, and educators, the research reveals persistent West-centric biases in how AI systems visualize non-Western subjects. The study found that basic prompts consistently generated stereotypical representations, with 82% of experts identifying poverty as the dominant narrative frame and images scoring an average of 3.2/10 for representational accuracy. While context-rich prompts achieved higher accuracy scores (7.8/10), the need for such detailed intervention highlights underlying systemic biases. The paper explores how these biases stem from unbalanced training datasets, Western-centric development practices, and the exclusion of indigenous knowledge systems. Drawing on a year-long experimental study, the research demonstrates how careful prompt engineering can improve cultural representation but argues that fundamental changes in AI development, including meaningful integration of diverse epistemologies and community partnerships, are necessary for truly inclusive systems. The findings have significant implications for how AI shapes cultural narratives and media representation in an increasingly algorithm-mediated world.

The child's eyes stared back at me through a screen, surrounded by what the AI had determined was an inevitable context: tattered clothing, a dusty street backdrop, and other unmistakable markers of poverty that Midjourney seemed to automatically attach to any prompt containing the words "Indian child." As I sat in my home in New Delhi, surrounded by the evidence of India's middle class reality, the contrast between my lived experience and the AI's perception felt like a punch to the gut. I had been attempting to generate images for a conference about science, technology, and society happening in New Delhi—conversations about gene editing, geopolitics, dematerialisation and agri-technology among others. Instead, each prompt seemed to channel decades of Western photojournalistic tradition that had consistently framed Indian childhood through the lens of hardship.

"An Indian child playing with friends," I typed, adding specific markers: "modern housing complex," "contemporary clothes," "middle-class setting." The result remained stubbornly stereotypical: children in worn clothing playing with makeshift toys in a dirt-covered setting. I adjusted the prompt: "Indian child in an international school classroom." Still, the AI rendered signs of deprivation—worn-out uniforms, sparse classroom settings, broken furniture. After dozens of attempts, the pattern became undeniable. The AI had been trained on decades of images that had prioritized a singular narrative of Indian childhood, and it was faithfully reproducing this limited perspective.

This wasn't merely frustrating—it was a stark illustration of how AI systems can fundamentally misunderstand and misrepresent non-Western experiences, reducing complex societies to their most stereotypical representations. This experience launches a deeper investigation into how AI image generation tools, particularly Midjourney, exhibit systemic biases when representing South Asian subjects. Through learnings from a year-long project dealing with over 1,000 AI-generated images, this paper reveals how these biases manifest and explores their broader implications for cultural representation in the digital age.



Prompt: a child in India



Prompt: An Indian child playing with friends, modern housing complex, contemporary clothes, middle-class setting.



Prompt: Indian child in an international school classroom

Scope of AI Bias

The Current State of AI Image Generation

Recent advancements in AI image generation have revolutionized digital content creation, with models like DALL-E, Stable Diffusion, and Midjourney processing millions of image requests daily (Phoenix & Taylor). These systems employ diffusion models and transformer architectures to generate images from textual descriptions, representing a significant leap in machine learning capabilities (Donvir et al., 2024). However, research indicates that these systems predominantly reflect Western artistic conventions and cultural perspectives, stemming from training datasets that are overwhelmingly sourced from Western digital archives and art collections (de Almeida & Rafael, 2024).

Overview of Midjourney's Capabilities and Limitations

Midjourney, as a leading AI image generation platform, demonstrates remarkable technical capabilities in artistic rendering and compositional coherence. While the platforms exhibit strong technical abilities in artistic rendering and handling complex aesthetic instructions, they show significant limitations in cultural contextualization. This limitation is part of a broader issue where AI systems' outputs are directly impacted by the quality and type of data they are trained on, leading to concerns about algorithmic bias (Arora et al., 2023).

Understanding Prompt Mechanics

The process by which AI systems interpret and respond to prompts involves multiple sophisticated computational mechanisms working in concert. At its foundation, AI systems process prompts through neural networks inspired by human brain architecture, employing multiple layers of interconnected nodes that progressively refine the system's understanding. This processing begins with tokenization and vectorization, where human language is converted into numerical representations that the AI can analyze. Critical to this process are attention mechanisms that help the system prioritize relevant information within the prompt, enabling more focused and contextual responses.

The AI's ability to generate responses stems from both its training data - comprising millions of data points that inform pattern recognition and contextual understanding - and its generative capabilities that allow for novel combinations of learned patterns. However, this process is inherently influenced by the biases present in training data, making ethical considerations and bias mitigation crucial aspects of prompt engineering. The entire mechanism operates within an adaptive framework where some systems can refine their response patterns over time through feedback loops, though this capability varies among different AI implementations. (Khan)

Generative AI Technical Architecture and Processing Mechanisms

Generative AI systems employ a structured implementation process with multiple distinct phases to transform inputs into desired outputs. The process begins with clearly defining the problem and desired outcomes, followed by data collection and preprocessing using appropriate tools and datasets. For example, this may involve gathering data through cameras, microphones, sensors, or existing curated datasets. The next critical phase is model selection, where an appropriate architecture (like VAEs, GANs, transformers, or diffusion models) is chosen based on the specific task requirements. The training phase then involves using substantial computational resources, often leveraging specialized hardware like GPUs or TPUs, to optimize the model parameters. During training, the model learns underlying patterns and statistical relationships from the training data. The model is then evaluated using various metrics to assess performance and quality. A fine-tuning phase follows, where hyperparameters are adjusted to optimize performance.

Once satisfactory results are achieved, the model is deployed for actual use. The final phase involves continuous monitoring and maintenance to ensure consistent performance and address any issues that arise. This systematic process reflects the intricate interplay between hardware requirements, software frameworks, and user experience considerations necessary for effective generative AI implementation. (Bandi et al., 2023)

Impact on Cultural Narratives and Media Representation

The increasing integration of artificial intelligence (AI) into the cultural landscape has profound implications, including its potential influence on religious expressions. Israeli historian Yuval Noah Harari predicts that AI could one day compose its own religious texts, attract followers, and even catalyze the creation of entirely new religions. This development signals not the end of history but the dawn of a new era, one where human culture—shaped over millennia—is reinterpreted and transformed by AI systems (Richmond, 2015). AI's ability to assimilate and regenerate cultural heritage introduces unique creative possibilities but also challenges the authenticity and agency of human narratives. Historically, human culture has served as the lens through which we experience reality, influencing our beliefs, preferences, and behaviors. With AI's capacity to craft unprecedented narratives and cultural artifacts, the prospect of perceiving reality through the prism of a non-human intelligence raises critical questions about authenticity and control. (Komal et al.)

The influence of AI-generated imagery on cultural narratives and media representation has become increasingly significant, with potential long-term implications for societal perceptions and cultural identity. Generative AI technologies are creating a “digital feedback loop” where biased representations reinforce existing stereotypes and shape future content creation. These biases

stem from unbalanced datasets that overrepresent dominant cultural norms while marginalizing minority perspectives. For instance, studies reveal that prompts like “playing basketball” predominantly generate images of African American men, reflecting societal stereotypes embedded in training data. As such content circulates online, it feeds back into AI systems, further entrenching these biases. This loop not only distorts public perceptions of beauty, identity, and societal roles but also perpetuates inequalities by normalizing exclusionary representations. The challenge lies in breaking this cycle through diversified training datasets and algorithmic fairness, ensuring generative AI fosters inclusivity rather than amplifying systemic disparities. Without these interventions, the iterative nature of AI risks solidifying harmful stereotypes in both media and societal attitudes. (Vázquez & Garrido-Merchán, 2024)

Decolonial Theory Framework

Decolonial theory provides a critical framework for examining how colonial power structures persist in modern technological systems, particularly in artificial intelligence. This theoretical approach emphasizes the ongoing influence of colonial dynamics in knowledge production, cultural representation, and economic systems long after formal colonialism has ended. The framework encompasses three key perspectives: decentering Western epistemologies, incorporating alternative knowledge systems, and promoting critical engagement with existing technological paradigms. When applied to algorithmic systems, this analysis reveals concerning patterns of algorithmic oppression, where AI systems perpetuate historical biases and discriminatory practices. This is particularly evident in areas such as criminal justice, facial recognition, and surveillance technologies, where marginalized communities face disproportionate negative impacts. Such examples demonstrate how colonial power structures can be inadvertently encoded into and amplified by modern AI systems, highlighting the importance of decolonial perspectives in technological development (Mohamed et al., 2020).

Algorithmic Dispossession

Algorithmic dispossession manifests through the concentrated control and development of AI technologies in the Global North, leading to significant power imbalances in the global AI landscape. (Png, 2022) This centralization is evident in how the Global North disproportionately controls global data flows and influences the digital economy's direction. The exclusion of developing nations from meaningful participation in AI governance represents a critical challenge for global security and equity. It further marginalizes these populations, particularly impacting international development where AI technologies are often proposed as solutions for complex developmental scenarios, yet the benefits primarily accrue to developed economies while developing nations face increasing technological dependency. As documented in Garcia's (2019) analysis,

the Global South faces systematic underrepresentation in key international discussions about AI, with many regions of Africa, Asia, Latin America and the Caribbean notably absent from crucial policy debates. This exclusion occurs despite these nations being potentially among the most vulnerable to negative impacts of military AI applications and autonomous weapons systems. The paper highlights how this creates a troubling dynamic where developing countries risk becoming mere “data-reservoirs and testbeds” for AI technologies while lacking the technical expertise and infrastructure to protect their interests. Beyond just participating in discussions, these nations often lack the capacity to develop their own AI national plans or implement effective countermeasures against more technologically advanced powers. This situation threatens to exacerbate existing global inequalities, creating a new divide between “AI-ready” and “not-ready” nations, while leaving developing countries particularly vulnerable to what the paper terms “data-predation and cyber-colonization” by more powerful states. (Garcia, 2019)

This creates a self-reinforcing cycle where technological and productivity gaps between Global North and South continue to widen, further entrenching existing power imbalances in global AI development and deployment. (Mohamed et al., 2020).

Root Causes: Beyond Technical Limitations

The origins and implications of AI bias extend far beyond mere technical constraints or limitations in system design. As Ferrara (2023) demonstrates, even seemingly minor biases within AI systems can trigger unpredictable cascading effects that ripple through society, particularly in how different cultures are represented and interpreted. These small initial biases can amplify and compound over time, creating increasingly significant disparities in how AI systems process, understand, and generate content related to different cultural contexts. The impact becomes particularly pronounced when these systems are deployed at scale across various applications and domains.

The fundamental challenge lies in what Ofosu-Asare (2024) identifies as “cognitive imperialism” - a deeply embedded systematic privileging of Western epistemologies and thought patterns in AI development. This bias manifests not just in the final outputs of AI systems, but in their foundational architecture, training methodologies, and the very frameworks used to conceptualize artificial intelligence. Western ways of knowing, understanding, and processing information become encoded as the default, universal standard, while alternative epistemologies and cultural frameworks of knowledge are marginalized or excluded entirely from the development process. This creates a self-reinforcing cycle where AI systems increasingly reflect and perpetuate Western cognitive paradigms while failing to adequately represent or understand other cultural perspectives.

Training data biases represent a fundamental challenge in artificial intelligence systems, stemming from historical inequities and sampling issues in dataset collection. Research has shown that widely used training datasets overrepresent dominant cultural perspectives and demographic groups while systematically excluding or misrepresenting minorities (Liu, 2024). This creates a cyclical problem where AI systems trained on biased data perpetuate and amplify existing societal biases.

The predominance of Western-centric development practices in AI reflects deeper structural issues within the technology industry. The concentration of AI research and development in North American and European institutions has led to systems that implicitly encode Western epistemological frameworks and value systems (Mohamed et al., 2020). This manifests in everything from the choice of problems being solved to assumptions about user needs and behaviors.

The limited inclusion of indigenous knowledge systems represents a critical gap in AI development that extends beyond mere representation. Indigenous ways of knowing often offer sophisticated frameworks for understanding complex systems and human-environment relationships that could enhance AI systems. However, these knowledge systems are frequently dismissed or overlooked in favor of Western scientific paradigms, despite their potential to contribute valuable perspectives on sustainability, collective decision-making, and ethical governance. These knowledge systems encompass multifaceted approaches to understanding phenomena, incorporating intergenerational wisdom, holistic perspectives, and relational ways of knowing that could fundamentally reshape how AI systems process and interpret information. For instance, indigenous perspectives on environmental stewardship often integrate long-term ecological observations with cultural practices and social responsibilities - a comprehensive approach that could inform more nuanced AI systems for environmental monitoring and resource management. Similarly, indigenous approaches to consensus-building and community-based decision-making could offer valuable insights for developing more equitable and culturally sensitive AI governance frameworks. The systematic exclusion of these knowledge systems not only perpetuates historical marginalization but also deprives AI development of rich epistemological traditions that could help address current limitations in machine learning approaches, particularly in areas such as contextual understanding, relational reasoning, and ethical decision-making. Moreover, indigenous knowledge systems often emphasize the interconnectedness of different domains of knowledge - a perspective that could help bridge the current gaps between AI's technical capabilities and its broader societal implications.

Methodology: Quantifying Cultural Bias

To move beyond anecdotal evidence, I developed a systematic evaluation framework. I assembled a diverse panel of twelve experts: three media representation scholars, four photographers who extensively covered Indian society, an international photographer, a sociologist specializing in class and representation in South Asia, and three educators from different socioeconomic contexts in India.

The evaluation process was designed to eliminate potential biases. Each expert received 30 pairs of images: one generated using basic prompts (“Indian child playing,” “Indian child at school”) and another using context-rich prompts (“A bright, modern classroom in an international school in New Delhi, India. Natural light streams through large windows. Smart board on wall, ergonomic student desks arranged in collaborative pods. Students wearing contemporary school uniforms. MacBooks and tablets on desks. Combination of Indian and international educational posters on walls. Clean, well-maintained space with modern LED lighting, central AC vents visible. Contemporary architectural details, vibrant blue and white color scheme, high ceilings”). The images were randomized, coded, and stripped of their prompting information. Experts didn’t know which images came from which prompting strategy, nor were they aware of the specific focus on poverty stereotypes in the study.

The evaluation form I developed asked experts to rate each image across multiple dimensions:

1. Representational Accuracy (0-10):

- Setting authenticity
- Clothing and appearance
- Activity contextuality
- Social class markers

2. Stereotype Presence (0-10):

- Poverty indicators
- Environmental markers
- Activity stereotyping
- Social context accuracy

The results were striking. Images generated with basic prompts scored an average of 3.2 out of 10 for representational accuracy, with 82% of experts identifying “poverty as the dominant narrative frame.” One expert, a veteran photographer who had documented Indian society for three decades, noted: “These images feel like they’re frozen in a Western photographer’s poverty porn portfolio from the 1980s.”

In contrast, images generated using context-rich prompts scored significantly higher, averaging 7.8 out of 10 for representational accuracy. However, the need for such detailed prompting revealed its own form of bias.

Another expert, a primary school teacher in an international school in India, observed: “I noticed in a lot of images children were rarely shown in active learning poses or engaged with technology. They were predominantly shown in passive or struggle-focused scenarios. This subtly reinforces the narrative of deprivation over development.”

Solutions and Future Directions

The implications of AI-generated imagery extend far beyond academic interest, revealing how these systems can amplify and institutionalize cultural biases at unprecedented scale and speed. As AI increasingly influences media, advertising, and public perception, systematic misrepresentation can reinforce harmful stereotypes and contribute to cultural erasure in ways both more subtle and pervasive than traditional media. When AI systems consistently generate images of Indian children in poverty-focused contexts, or depict South Asian settings through an outdated colonial lens, they don't merely reflect existing biases – they encode and perpetuate them through seemingly objective technological systems. As Broussard (2024) notes, these technological biases have real-world consequences for how societies understand and value diverse cultural expressions.

The impact manifests across multiple domains: in journalism, where generated images might illustrate stories with stereotypical visuals; in advertising, where algorithms perpetuate limited cultural narratives; in educational materials, where generated content presents students with skewed representations; and in social media, where these images shape public discourse and personal perceptions. The automated nature of these systems means they can generate and distribute biased representations at scale, potentially overwhelming more nuanced human-created depictions. This creates a feedback loop where AI-generated stereotypes influence human perceptions, which inform future training data, further entrenching these biases. The economic implications are equally concerning – when AI consistently misrepresents certain communities, it affects everything from market research to product development, potentially leading to systematic economic exclusion. The psychological impact on misrepresented communities is significant as they encounter distorted versions of their cultural identity in digital spaces, particularly affecting younger generations who grow up in an environment where generated content plays an increasingly central role in shaping cultural narratives.

Research on prompt engineering has identified several effective strategies for improving AI model outputs across cultural contexts. Studies have demonstrated

that carefully constructed prompts incorporating cultural context, explicit fairness considerations, and diverse perspectives can help mitigate inherent model biases (Khan). Particularly successful approaches include using counterfactual examples, implementing systematic bias-checking frameworks, and employing culturally-informed evaluation criteria.

Indigenous epistemologies offer transformative potential for reconceptualizing AI development paradigms. Drawing from Mohamed et al.'s (2020) analysis, this includes establishing a critical technical practice that questions power imbalances and implicit value systems; implementing reverse tutelage where marginalized communities actively shape AI development; and fostering new forms of affective and political community beyond paternalistic approaches. The authors advocate for concrete steps such as diversifying AI teams, auditing training datasets for cultural biases, expanding evaluation metrics, and implementing cultural verification systems. They emphasize that technical solutions alone cannot create meaningful progress – fundamental changes are needed in how AI systems are conceptualized and developed, including the integration of indigenous knowledge systems, co-development strategies with affected communities, and mechanisms for meaningful intercultural dialogue.

Drawing on Liu et al.'s (2024) research, culturally adaptive AI requires synergy between technological advancement and authentic community partnership. Their framework reveals that meaningful representation encompasses shared knowledge, value systems, and societal norms. This insight proves particularly salient when examining MidJourney's West-centric limitations, underscoring the necessity for AI systems to progress beyond surface-level recognition toward genuine cultural embodiment. The authors' distinction between "deep" and "surface" adaptation illuminates how meaningful change demands collaborative development with communities from the outset, rather than retrospective consultation. When applied to visual generation platforms, this paradigm suggests that enhancing South Asian representation demands a fundamental reimagining of how these systems understand and encode cultural perspectives, developed in partnership with South Asian voices.

The investigation into cultural alignment of Large Language Models reveals critical insights about the intersection of AI systems and human cultural diversity. Through examination of model responses across languages, demographic dimensions, and cultural contexts, this study demonstrates that current LMs exhibit significant disparities in their cultural awareness and representation. The introduction of Anthropological Prompting represents a promising step toward more culturally nuanced AI systems, but considerable work remains to achieve truly inclusive and culturally competent language models (AlKhamissi et al., 2024). Future research must focus on expanding datasets, refining technical approaches, and deepening collaboration with diverse communities to ensure cultural adaptation efforts authentically serve and represent the full spectrum of human experience.

My year long experimentation with generating images on Midjourney tells a story of slow but meaningful progress: where my initial attempts at generating images of future technologies, scientific concepts and current conversations in an Indian context yielded stereotypical representations most of the time, refined prompts now achieve cultural accuracy more often than not. Through careful prompt engineering and systematic documentation of effective approaches, I can consistently generate images that respect and accurately represent South Asian cultural elements. Yet the need for such extensive intervention highlights the work still needed to create truly inclusive AI systems.

This improvement didn't emerge from technical adjustments alone; it came from the integration of local knowledge, from understanding contemporary South Asian landscapes, from knowing what a modern Indian classroom really looks like, and from living these realities daily.

The journey from frustration to incremental progress reveals a broader truth about artificial intelligence: these systems are mirrors, reflecting not just our world but our worldview. When that worldview is limited, so are the images it produces. As AI image generation increasingly shapes how we visualize and understand each other, we can't afford to let outdated narratives calcify into algorithmic bias. Each stereotypical image generated doesn't just misrepresent a moment; it reinforces a perspective that ripples through media, education, and cultural understanding.

"The real question isn't about technology anymore," reflects Dr. Shankar, one of the specialists who evaluated tens of AI-generated images for my study. "It's about whose stories get to be told, and who gets to tell them." As millions of images are generated daily, shaping perceptions and reinforcing narratives, her words carry a particular urgency. The technology exists to do better. The expertise exists to guide us. The question that remains is whether we have the will to listen to the voices that have been left out of the conversation for too long.



Prompt: A bright, modern classroom in an international school in New Delhi, India. Natural light streams through large windows. Smart board on wall, ergonomic student desks arranged in collaborative pods. Students wearing contemporary school uniforms. MacBooks and tablets on desks. Combination of Indian and international educational posters on walls. Clean, well-maintained space with modern LED lighting, central AC vents visible. Contemporary architectural details, vibrant blue and white color scheme, high ceilings

References

- Phoenix, J., & Taylor, M. (n.d.). Prompt engineering for generative AI. O'Reilly Online Learning. <https://learning.oreilly.com/library/view/prompt-engineering-for/9781098153427/>
- Donvir, A., Panyam, S., Paliwal, G., & Gujar, P. (2024b). The Role of Generative AI Tools in Application Development: A Comprehensive Review of Current Technologies and Practices. 2024 International Conference on Engineering Management of Communication and Technology (EMCTECH), 1-9. <https://doi.org/10.1109/emctech63049.2024.10741797>
- de Almeida, F., & Rafael, S. (2024). Bias by default.: Neocolonial visual vocabularies in AI image generating design practices. Extended Abstracts of the CHI Conference on Human Factors in Computing Systems, 1-8. <https://doi.org/10.1145/3613905.3644053>
- Arora, A., Barrett, M., Lee, E., Oborn, E., & Prince, K. (2023). Risk and the future of AI: Algorithmic bias, Data Colonialism, and marginalization. *Information and Organization*, 33③, 100478. <https://doi.org/10.1016/j.infoandorg.2023.100478>
- Khan, I. (n.d.). The Quick Guide to Prompt Engineering. O'Reilly Online Learning. <https://learning.oreilly.com/library/view/the-quick-guide/9781394243327/c02.xhtml#head-2-5>
- Bandi, A., Adapa, P. V., & Kuchi, Y. E. (2023). The power of Generative AI: A review of requirements, models, input-output formats, evaluation metrics, and challenges. *Future Internet*, 15⑧, 260. <https://doi.org/10.3390/fi15080260>
- Komal, Singh, A., Sethu, S., & Chaudhary, R. (n.d.). Artificial Intelligence in Storytelling: A Critical Analysis of Narrative Authenticity, Cultural Impact, and the Future of Creative Industries. *View of Artificial Intelligence in storytelling: A critical analysis of narrative authenticity, cultural impact, and the future of Creative Industries*. <https://bpasjournals.com/library-science/index.php/journal/article/view/1284/802>
- Png, M.-T. (2022). At the tensions of South and north: Critical roles of Global South stakeholders in AI Governance. 2022 ACM Conference on Fairness, Accountability, and Transparency, 1434-1445. <https://doi.org/10.1145/3531146.3533200>
- Richmond, S. (2015). *Superintelligence: Paths, dangers, strategies*. by Nick Bostrom. Oxford University Press, Oxford, 2014, pp. XVI+328. hardcover: \$29.95/ £18.99. ISBN: 9780199678112. *Philosophy*, 91①, 125-130. <https://doi.org/10.1017/s0031819115000340>
- Vázquez, A., & Garrido-Merchán, E. C. (2024). A Taxonomy of the Biases of the Images Created by Generative Artificial Intelligence.
- García, E. (2019). The militarization of Artificial Intelligence: A wake-up call for the Global South. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.3452323>
- Ferrara, E. (2023). The butterfly effect in artificial intelligence systems: Implications for AI bias and fairness. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.4614234>
- Ofori-Asare, Y. (2024b, September 19). Cognitive imperialism in artificial intelligence: Counteracting bias with indigenous epistemologies - ai & society. *SpringerLink*. <https://link.springer.com/article/10.1007/s00146-024-02065-0>
- Liu, Z. (2024). Cultural bias in large language models: A comprehensive analysis and mitigation strategies. *Journal of Transcultural Communication*. <https://doi.org/10.1515/jtc-2023-0019>
- Liu, C. C., Gurevych, I., & Korhonen, A. (2024, June 6). Culturally aware and adapted NLP: A taxonomy and a survey of the state of the art. *arXiv.org*. <https://arxiv.org/abs/2406.03930>
- AlKhamissi, B., ElNokrashy, M., Alkhamissi, M., & Diab, M. (2024). Investigating cultural alignment of large language models. *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 12404-12422. <https://doi.org/10.18653/v1/2024.acl-long.671>