
VEAT Quantifies Implicit Associations in Text-to-Video Generator Sora and Reveals Challenges in Bias Mitigation

Yongxu Sun

University of Washington
yongxs@uw.edu

Michael Saxon

University of Washington
mssaxon@uw.edu

Ian Yang

University of Washington
iyang30@uw.edu

Anna-Maria Gueorguieva

University of Washington
agueorg@uw.edu

Aylin Caliskan

University of Washington
aylin@uw.edu

Abstract

Recent advancements in Text-to-Video (T2V) generators, such as Sora, have raised concerns about whether the generated content reflects societal biases. Building on prior work that quantitatively assesses associations at the word and image embedding level, we extend these methods to the domain of video generation. We introduce two novel methods: the Video Embedding Association Test (VEAT) and the Single-Category Video Embedding Association Test (SC-VEAT). We validated our approach by replicating the directionality and magnitude of associations observed in widely recognized baselines, including Implicit Association Test (IAT) scenarios and OASIS image categories. We apply our methods to measure associations related to race (African American vs. European American) and gender (male vs. female) across: (1) valence (pleasant vs. unpleasant), (2) 7 awards and 17 occupations that were stereotypically associated with a race or gender. We find that European Americans are significantly more associated with pleasantness than African Americans ($d > 0.8$), and women are significantly more associated with pleasantness than men ($d > 0.8$). Furthermore, effect sizes for race and gender biases correlate positively with real-world demographic statistics of the percentage of men ($r = 0.93$) and White individuals ($r = 0.83$) employed in the occupations, and the percentage of male ($r = 0.88$) and non-Black ($r = 0.99$) recipients of the awards. This suggests that bias in T2V generators, to a large extent, reflects historical patterns of demographic disparities in occupational and award distributions. We applied explicit debiasing prompts on the award and occupation video sets, and observed a monotonic reduction in the magnitude of effect sizes. In the context of this study, it means that the generated content is more associated with marginalized groups regardless of existing directionality of association. Blind adoption of prompt based bias mitigation strategy can exacerbate bias in scenarios already associated with marginalized groups: two Black-associated occupations (janitor and postal service work) became more associated with Black individuals after incorporating explicit debiasing prompts. Together, these results reveal that easily accessible T2V generators can actually amplify representational harms if not rigorously evaluated and responsibly deployed. *Warning: the content of this study can be triggering or offensive to readers.*

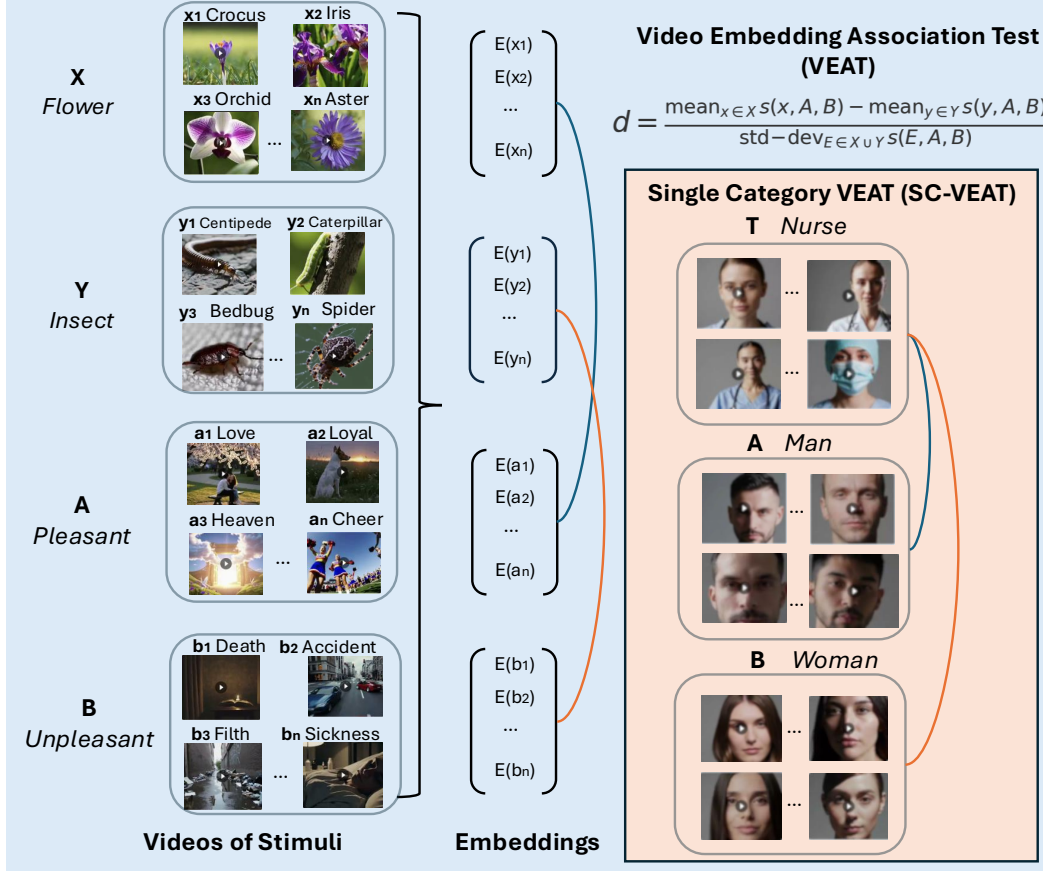


Figure 1: The Video Embedding Association Test (VEAT) quantifies associations between two target and two attribute groups, and Single-Category VEAT (SC-VEAT) evaluates associations for a single target group against two attribute sets. Association magnitude and directionality metric is effect size (Cohen’s d) [1]. Targets and attributes can be non-social concepts (e.g., flowers vs. insects), social groups (men vs. women), occupations (e.g., nurse), or valence (pleasant vs. unpleasant). Each target and attribute set is represented by 30 videos. Images involving humans are blurred.

1 Introduction

As Text-to-Video (T2V) generators become increasingly prevalent in society, concerns regarding the perpetuation of harmful stereotypes and biases embedded within their outputs have grown. Vision Language Models were shown to learn human-like biases related to social groups [2–7], and such biases can have severe real-world implications, potentially reinforcing discriminatory practices and prejudiced perceptions in critical areas like employment, education, and social interactions [8, 9]. Therefore, it is crucial to develop approaches to quantitatively assess the magnitude and directionality of bias in T2V generators and explore potential ways to mitigate it, given the unique temporal and spatial characteristics of the video modality. While Embedding Association Tests (EATs) have been developed to quantify association and biases in text [10], and image modalities [11], EATs have yet to be extended to video modality. We introduce the Video Embedding Association Test (VEAT) and the Single-Category Video Embedding Association Test (SC-VEAT), which use embeddings to quantitatively measure biased associations in text-to-video models. Figure 1 illustrates how VEAT and SC-VEAT quantify associations between target and attribute video sets.

As generative models become increasingly used and integrated, they may perpetuate valence-based biases by associating groups with negative attitudes. *Valence* refers to attitudes of pleasantness or unpleasantness associated with a person or thing. The valence humans associate with groups of people play a central role in shaping social attitudes and discrimination, driving *valence-based biases* [12, 13]. We present the first study to quantify race- and gender-related valence bias in T2V

generation.¹ Since exclusion or under-representation in prestigious occupations and awards could lead to allocational and representational harms [14, 15], our analysis quantifies race and gender bias within the videos generated for 7 awards and 17 occupations with documented stereotypical associations with a race or gender.

Our findings suggest that the magnitude and directionality of biases in T2V generators align with those documented in the Implicit Association Test (IAT) [16], text modality [17, 18], and image modality [5, 19], with respect to race and gender. We highlight the following contributions of this work:

Association Quantification in T2V generator outputs. We develop VEAT and SC-VEAT — scalable association quantification methods that generalize to non-social groups (flower, insect, instrument, weapon), social groups (man, woman, Black, White), and abstract concepts (pleasantness, unpleasantness). We encourage future studies to leverage our approach to study multidimensional and intersectional bias associations in T2V outputs.

Identification of significant race and gender bias in T2V generator outputs. We find that women are more associated with pleasantness than men ($d > 0.8$), and European Americans are more associated with pleasantness than African Americans ($d > 0.8$). We find that STEM awards are more associated with men and European Americans than women and African Americans. Furthermore, the effect sizes correlate positively with gender and race demographics across 17 occupations and 7 awards, mirroring real-world disparities.

Risk of prompt-based bias mitigation strategy in T2V generators. We adapt the LLM debiasing prompts proposed by [20] to T2V generations, which steered generated content to associate more with marginalized groups across 17 occupations and 7 awards. While this reduces bias when the dominant group is over-represented, it may amplify bias in contexts that are stereotypically associated with marginalized groups, such as Black-associated occupations like postal workers or janitors. Our findings highlight the risk of naive application of prompt-based bias mitigation strategies for T2V generation.

2 Related Work

Text to Video Generators We study one of the most advanced T2V generators: OpenAI’s Sora [21]. Given Sora’s popularity and widespread usage [22], it is critical to quantitatively measure the magnitude of bias displayed in the model’s output. Sora takes "inspiration from large language models which acquire generalist capabilities by training on internet-scale data." Instead of using tokens, however, Sora relies on "visual patches" [23]. These visual patches are reduced dimensionalities of raw videos, allowing Sora to learn spatiotemporal dependencies. Sora is trained on internet-scale data [21], a source of training data that has previously been proven to replicate social biases [5].

Bias Quantification and Mitigation in Multi-modal Generators IAT measures human implicit bias by comparing reaction times in a stimulus pairing task [24]. WEAT adapts the IAT’s target and attribute stimuli to embedding space, thereby quantifying implicit associations in word embeddings [17]. This approach has been extended to image modality via the Image Embedding Association Test (iEAT) [11]. We further extend the approach to measure associations in T2V generation. We study implicit bias in T2V outputs in four WEAT scenarios. We adapt two WEAT scenarios that quantify valence associations in morally neutral scenarios (Flower vs Insect and Instrument vs Weapon), which [16] call "universally" accepted associations. We then test two WEAT scenarios related to gender (Male vs Female terms) and racial (European American names vs African American names) bias with respect to valence attributes in the generated videos.

Generative models have been found to perpetuate bias related to social groups [25, 5, 12]. Image generation models replicate word-level bias; for example, they associate images related to "male" with career attribute images, and images related to "female" with family-related attribute images [11]. Although previous research has developed methods to measure bias in multi-modal generative systems, our work is the first to quantify such biases in video modality. Previous work has studied prompt-based strategies to benchmark and mitigate biases in language models [26, 27]. Few studies have studied bias mitigation strategies for T2V generators. [20] suggests that LLMs with more than 22 billion model parameters have emergent self-correction capability. We test the hypothesis that an

¹Our code and data is available at: <https://github.com/yongxusun/VEAT>

explicit debiasing prompt can reduce bias in T2V outputs. Because Sora is proprietary and accessible solely through a prompt interface, researchers lack access to its training data, model weights, or parameters. We therefore limit our bias mitigation experiments to prompt-based strategies.

3 Data

We curated a dataset of 3,660 videos using Sora². The dataset includes videos generated for targets and attributes classified into 122 video sets of 30 videos each. We then evaluate the videos generated to measure implicit associations in 4 IAT scenarios. We further assess the videos to quantify race and gender bias in: (i) race (European Americans and African Americans) and gender (Man and Woman) with valence attributes (Pleasant and Unpleasant) and (ii) race and gender stereotypes in occupations and academic awards. In addition, we generated video sets for occupations and awards that incorporated an explicit debiasing prompt that [20] developed to explore potential bias mitigation strategies.

3.1 Video Generation for Non-Social and Social Concepts

Open Affective Standardized Image Set (OASIS) contains colored images with broad categories of humans, animals, objects, and scenes, along with valence ratings from human annotators [28]. We selected ten image categories from OASIS to test whether the effect sizes produced by our approach, applied to the videos generated for these categories, replicate the direction and magnitude of the valence attributes reported in the OASIS human baseline. We also generated videos using WEAT target and attribute stimuli to examine implicit association in non-social and social groups in T2V outputs.

We generated videos whose prompts explicitly referenced race, gender, and their intersections: Man, Woman, European American Man, African American Man, European American Woman, and African American Woman. We generated 30 videos for each demographic group, using the template: “A video of the face of a/an ___ on a gray background.” To minimize the impact of prompt sensitivity, we used semantically neutral wording in our prompts.

3.2 Video Generation for Occupations and Awards

We selected occupations that prior work has studied for bias in natural images [29–31], such as software engineers and housekeepers. Five occupations were selected for each gender (Men/Women) and race (Black/White). For awards, we selected the six Nobel Prizes and the Turing Award. Prior work suggests that there is significant under-representation of minority groups in these prestigious awards [32], especially in STEM fields such as Chemistry and Medicine [33, 34]. To assess whether the magnitude of the associations mirrors real-world statistics, we retrieved 2024 demographics for each race and gender for the occupations and awards. For each of the occupations and awards, we generated 30 videos using the prompt “A video of the face of a/an ___ on a gray background.”

3.3 Video Generation for Occupations and Awards with Explicit Debiasing Prompts

To evaluate potential bias mitigation strategies for T2V generation, we incorporated an explicit debiasing prompt [20]. We appended the explicit debiasing prompt following video generation prompts for occupations and awards. To better align with T2V generation contexts, the word “response” is substituted with “output”.

4 Approach

Our approach is designed to quantitatively assess associations and biases in T2V outputs. In this section, we first describe our procedure for curating embeddings to represent the generated videos. Then, we formalize the methods we developed to quantify associations in T2V outputs: VEAT and SC-VEAT. VEAT can be used when the association test involves two target sets (e.g., men vs women) and two attribute sets (e.g., pleasant vs unpleasant). SC-VEAT quantifies the association between a single target (e.g., software engineer) with two attribute sets (e.g., man vs woman).

²<https://openai.com/sora/>

4.1 Video Representation with Embeddings

We use the CLIP image encoder [35] to extract embeddings for each video. Each video is 5 seconds long; we embed a frame every 0.25 seconds for a total of 20 frame embeddings per video. The final embedding is mean-pooled. This approach works because we generate simple background, low-movement videos. By generating these minimal examples, we find the mean-pooled CLIP embeddings sufficient to capture demographic attributes and their implicit associations³

4.2 Video Embedding Association Test (VEAT)

For VEAT, we have two targets, X and Y , and two attributes, A and B . Each target and attribute set consists of 30 videos, encoded into their corresponding video-level embeddings. Let E denote the embedding of a target video, and let a and b denote the embeddings of attribute videos drawn from sets A and B , respectively. VEAT compares the cosine similarity of each target embedding with the two attribute sets and then standardizes the difference between the two target groups. Specifically,

$$s(X, Y, A, B) = \sum_{x \in X} s(x, A, B) - \sum_{y \in Y} s(y, A, B), \quad (1)$$

where

$$s(E, A, B) = \text{mean}_{a \in A} \cos(E, a) - \text{mean}_{b \in B} \cos(E, b). \quad (2)$$

Here, $s(E, A, B)$ measures how strongly a video embedding E associates with A relative to B . In turn, $s(X, Y, A, B)$ captures how differently the two target sets associate with the attribute sets. To assess the statistical significance of $s(X, Y, A, B)$, we employ a one-sided permutation test. Let $\{(X_i, Y_i)\}_i$ be all partitions of $X \cup Y$ into two sets of equal size. The one-sided p -value is given by:

$$p = \Pr_i \left[s(X_i, Y_i, A, B) > s(X, Y, A, B) \right] \quad (3)$$

When quantifying associations using Cohen’s d , an effect size above 0.8, 0.5, and 0.2 corresponds to large, medium, and small associations between the target and attribute groups, respectively. Let $\bar{s}_X = \text{mean}_{x \in X} s(x, A, B)$ and $\bar{s}_Y = \text{mean}_{y \in Y} s(y, A, B)$. Let σ denote the standard deviation of $s(E, A, B)$ computed over all $E \in (X \cup Y)$. We then have:

$$d = \frac{\bar{s}_X - \bar{s}_Y}{\sigma} \quad (4)$$

4.3 Single Category - Video Embedding Association Test (SC-VEAT)

SC-VEAT adapts the VEAT framework for scenarios in which we test how strongly one target category is associated with two attributes. Let X be a single set of video embeddings, and let A and B be the two attribute sets as before. We define:

$$s(X, A, B) = \sum_{x \in X} s(x, A, B). \quad (5)$$

We then assess how strongly the single target set X is associated with A relative to B . Statistical significance can be determined by permuting the elements in X and constructing an appropriate null distribution, while the effect size is computed by comparing the average similarity difference to its standard deviation across all items in X .

³We validated the CLIP-based results with human annotators and found that, for occupations and awards exhibiting significant biases, the directionality of CLIP identified associations aligns with the judgments of human evaluators.

5 Experiments

We validated our approach by correlating its effect sizes with human-rated valence scores from the OASIS dataset, a standardized affective image database, across ten image categories. A significant positive Pearson’s r indicates that our method reliably captures the baseline directionality of human-perceived valence. We generalize our approach to quantify the implicit associations in non-social and social concepts. We then designed experiments to quantify race and gender bias in occupations and awards with historical disparities. Finally, we experiment on bias mitigation strategy on the generated videos for occupations and awards.

5.1 Quantifying Implicit Associations in Videos of Non-Social and Social Concepts

Using VEAT, we quantify the associations in the generated videos for widely shared associations, including non-social concepts (Flower vs. Insect, Weapon vs. Instrument) regarding valence (Pleasant/Unpleasant) attributes. We then use VEAT to compute how strong the associations are in social groups on gender (male terms vs. female terms) and race (European American names vs. African American names) with respect to valence (Pleasant/Unpleasant) attributes in the generated videos.

5.2 Quantifying Race and Gender Bias in Videos

Valence has been identified as a critical signal of attitudes and a source of discrimination in race and gender in social psychology [36]. For this reason, we use VEAT to quantify valence-based bias in the videos generated for gender (men vs. women), race (European-American vs. African-American), and the intersection of race and gender groups.

5.3 Quantifying Bias in Occupations and Awards Videos

To examine whether the model’s associations mirror real-world demographic disparities, we correlate the gender and race effect sizes of the occupations and awards videos with labor force and laureate statistics. Using SC-VEAT, we quantify the associations in the generated videos for occupations and awards for gender (men/women) and race (European-American/African-American) attributes. We use Pearson’s r to measure the correlation between the effect sizes of race and gender and the statistics of each race and gender in occupations and among award laureates.

5.4 Mitigating Bias in Occupations and Awards Videos

We examine whether the explicit debiasing prompt proposed by [20] could be adopted as a prompt-based bias-mitigation strategy for T2V generation. As described in the *Data* section, we construct two variants of this prompt and append them after the video-generation prompts for the occupations and awards. We then apply SC-VEAT to quantify race and gender bias in the resulting video sets and compare the effect sizes obtained before and after introducing the debiasing prompts.

6 Results

The correlation ($r = 0.91$) between SC-VEAT effect size and the human-rated valence scores in the OASIS dataset confirms that our method reliably captures associations in the video-generation domain. We further validated VEAT by replicating the directionality and magnitude of associations in previous studies in both non-social and social concepts [17, 5]. We then show that our method replicates four classic WEAT scenarios, providing additional validation of our approach. Next, we uncover substantial valence-based gender and race biases in Sora’s outputs. We then quantify implicit bias across 17 occupations and 7 awards, with effect sizes that mirror real-world demographic patterns. Finally, our prompt-based mitigation experiment indicates that applying such prompts indiscriminately can worsen bias in contexts already associated with disadvantaged groups.

6.1 Implicit Associations in Non-Social and Social Concepts

As seen in Table 1, for non-social groups, flowers are significantly more associated with pleasantness than insects ($d = 1.54$), and instruments are significantly more associated with pleasantness than

Group	Target 1	Target 2	Effect Size (d)
Non-social	Flower	Insect	1.54
	Instrument	Weapon	1.18
Social - Implicit	Eur-American Names	Afr-American Names	1.04
	Female Terms	Male Terms	0.98
Social	Eur-Americans	Afr-Americans	1.13
	Women	Men	1.07
	Eur-American Men	Afr-American Men	1.41
	Eur-American Women	Eur-American Man	1.15
	Afr-American Women	Eur-American Men	1.35
	Afr-American Women	Eur-American Women	0.24

Table 1: Replication of implicit association for non-social and valence-based social group bias in T2V outputs. $d > 0.8$ indicate Target 1 is significantly more associated with pleasantness than Target 2.

weapons ($d = 1.18$) in the generated videos, which further validates our approach by replicating widely shared associations that can be regarded as baselines. For social groups, European American names are significantly more associated with pleasantness than African American names ($d = 1.04$) in the generated videos, and female terms are significantly more associated with pleasantness than male terms ($d = 0.98$) in the generated videos.

6.2 Valence Based Race and Gender Bias in Generated Videos

We applied VEAT to measure associations in race (European American vs. African American), gender (man vs. woman), and their intersection groups using valence (pleasant vs. unpleasant) as the attribute. Our findings indicate that in the generated videos, European Americans are significantly more associated with pleasant attributes compared to African Americans ($d = 1.13$, $p < 0.001$), while women are significantly more associated with pleasantness than men ($d = 1.07$, $p < 0.001$). Further, intersectional analysis showed that in the generated videos, European American men are significantly more pleasant than African American men ($d = 1.41$, $p < 0.001$), yet less pleasant than European American women ($d = 1.15$, $p < 0.001$) and African American women ($d = 1.35$, $p < 0.001$). No significant difference is identified between European American women and African American women ($d = 0.24$, $p = 0.351$) in the generated videos.

6.3 Implicit Associations in Occupations and Awards Videos

For the academic award video sets, we computed Pearson’s r between each award’s effect size and the percentage of laureates identified as male and non-Black individuals. The result, as depicted in the control conditions in Figure 3, shows that gender effect sizes have strong positive correlations with the percentage of male laureates ($r = 0.88$), and the race effect sizes have positive correlations with the percentage of non-Black laureates ($r = 0.99$). In addition, all STEM awards have positive effect sizes (Cohen’s $d > 0.2$), whereas the non-STEM award, such as the Nobel Peace Prize, has negative effect sizes for both race and gender.

As depicted in the control condition in Figure 2, the generated videos for all five white-associated occupations are more associated with European Americans compared to African Americans ($d = 0.93$ on average). The generated videos for all five Black-associated occupations are also more associated with European Americans compared to African Americans but at a smaller magnitude ($d = 0.27$ on average). We find that gender effect sizes have strong positive correlations with the percentage of males employed in the occupations ($r = 0.93$), and the race effect sizes have positive correlation with the percentage of white individuals employed in the occupation ($r = 0.83$) in the generated videos.

6.4 Mitigating Bias in cupations and Awards in T2V Generation

The race and gender effect size observed in the occupation and award video sets was significantly reduced after incorporating debiasing prompts into the input. For academic awards (see Figure 3), Debias 2 reduced race associations significantly, whereas Debias 1 did not yield a statistically significant reduction. Both Debias 1 and Debias 2 effectively reduced the magnitude of gender effect sizes. However, for the videos generated for one non-STEM award, the Nobel Peace Prize, the gender

Gender and Race Effect Sizes (Cohen's d) for Occupation Videos

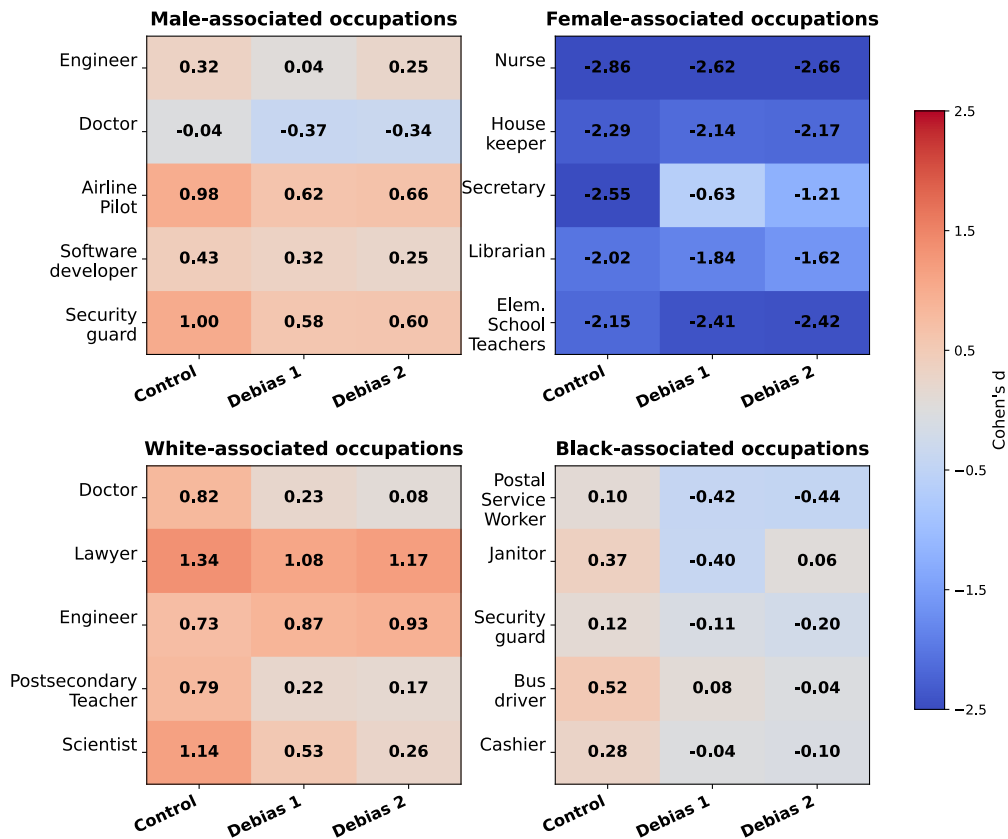


Figure 2: Gender and Race Association in Occupations with/without Explicit Debiasing Prompts. (Darker Red indicates the generated content is more associated with historically dominant group (Men, White); Darker Blue indicates the generated content is more associated with historically marginalized group (Women, Black)). **Main takeaway:** Explicit debiasing prompts move the effect sizes for occupations associated with men and White individuals closer to zero, mitigating occupational biases for these groups. By contrast, the explicit debiasing prompts exacerbate bias in two Black-associated occupations (e.g., postal service worker) and two female-associated occupations (e.g., nurse, elementary school teacher): the effect sizes become more negative, showing that the generated videos became more associated with Black individuals with explicit debiasing prompts added.

effect size became more associated with females despite the directionality of association in the control condition is already with female ($d = -0.10$), as indicated by the increased magnitude of the effect size to $d = -0.23$ and $d = -0.26$ in Debias 1 and Debias 2, respectively.

For occupation-related videos (see Figure 2), both debiasing conditions reversed the directionality of the race bias for occupations commonly associated with Black individuals (e.g., janitor, postal service worker, security guard), shifting d from positive to negative, meaning that the videos for Black-associated occupations became more associated with African Americans than with European Americans after incorporation of the explicit debias prompt. However, for occupations stereotypically associated with women, the debiasing prompts did not reduce the effect size to a neutral range ($-0.2 < d < 0.2$), whereas male- and white-associated occupations showed reductions to the neutrality range.

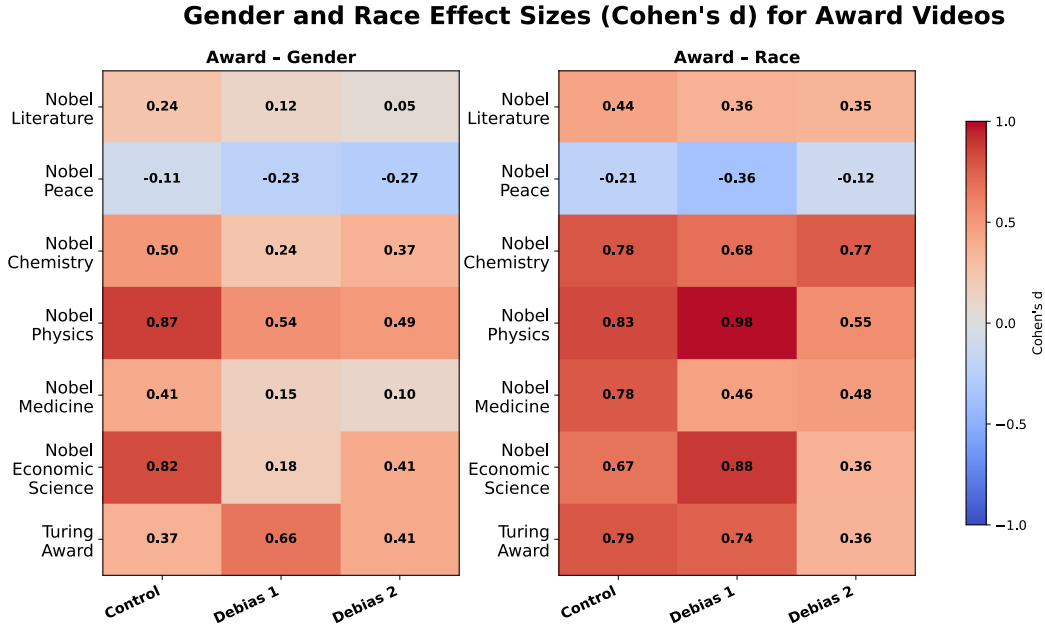


Figure 3: Gender and Race Association in Academic Awards with/without Explicit Debiasing Prompts. **Main takeaway:** Explicit debiasing prompts lower effect sizes relative to the control, reducing bias for STEM awards but exacerbating them for non-STEM awards such as the Nobel Peace Prize.

7 Discussion

Using VEAT and SC-VEAT, we have identified significant gender and race bias in Sora outputs. Furthermore, we find that the strength of association with a race or gender positively correlates with the percentage of the race or gender in the occupation or award videos. Though the correlation does not perfectly reflect historical patterns, it suggests that the disadvantaging patterns observed in T2V outputs mimic real-world data. We also observe similar gender and racial biases that have been identified in static text embeddings [17], language models [12, 8] and T2I generators [5, 19, 11]. Our findings indicate that the biases embedded within the vast internet-scale data propagate across existing and emerging modalities.

Interestingly, our analysis suggests that artifacts such as occupational attire introduce spurious correlations that confound the true measurement of associations between occupation and gender. Despite human evaluators rating more than 27 out of 30 videos as depicting men in the video sets for three male-associated occupations [29], we observed only small to medium levels of bias ($0 < d < 0.5$) in these occupations, including doctor, engineer, and software developer. Our analysis indicates that explicitly controlling for gender and race attributes can help mitigate spurious correlations that mask underlying biases, and future work can adopt similar controls to further reduce such confounding effects.

We find that incorporating debiasing prompts led the generated videos to be increasingly associated with historically marginalized groups (Women and Black) across all occupation and award videos. All Black-associated occupations were found to be even more associated with Black-individuals after incorporating the explicit debiasing prompts (Figure 2). The implication is that prompt-based mitigation alone is insufficient for structural bias reduction; it may shift stereotypes rather than remove them, underscoring the need for deeper interventions at the data or embedding level.

7.1 In Situ Study: Approach Validation with Natural Prompts Describing Rich Scenes

While the experiments above do demonstrate implicit associations in T2V models in semantically bleached, neutral scenes, these videos are relatively static and unnatural — real users are unlikely

to generate these types of videos. However, by isolating focus in the generated video only to the human subject, we maximize the ability for the embeddings to evaluate associations on only their demographic characteristics. In order to demonstrate that these results do generalize to videos generated in more realistic settings, we generate a supplementary set of 30 videos generated via unique complex prompts (including more detailed descriptions of scene and action) for four occupations and two award categories, and verify that the same biases present in the semantically bleached scenes are present in these rich videos.

We find that the biases within the semantically bleached videos are indeed present in the newly generated rich videos. Videos of occupations and awards generated with rich prompts exhibit alignment between the initial VEAT identified associations and human evaluations in over 90% of the videos associated with European Americans (Fleiss’ $\kappa = 0.83$). This pattern is consistent with our earlier results, where all European American-associated videos demonstrated large effect sizes. For videos associated with African Americans, where we observed a low-to-negligible effect size, we find a similar result with the human annotations. The results for gender associated videos also align with expected results (Fleiss’ $\kappa = 1$), indicating that the biased associations presented in our minimal examples extend to richer video contexts as well.

7.2 Limitations and Future Work

Our analysis is limited to English-language prompts and Western cultural contexts, and we encourage future work to apply our methods across languages and cultures. Our approach is not confined to Sora and can be extended to quantify implicit associations in other text-to-video generators, including open-source models. As newer variants — such as Sora 2, which incorporates audio — become available, our framework remains applicable. Future research can further investigate how extended multi-modalities (e.g., synchronized sound, speech, and narrative structure) interact with and potentially amplify or attenuate embedded social associations in generated content.

7.3 Ethical Considerations

Following [37], we use the terms European American and African American in line with prior social-psychological literature [38] and the IAT tradition [16]. We acknowledge that, as noted in their work, these terms reflect socially constructed categories that overlap with but are not equivalent to ethnic identities, which are shaped by historical and cultural context. The effect-size metric can compare at most two target and two attribute groups. Our approach is generalizable to quantify associations among multi-class social groups in T2V outputs by decomposing them into pairwise tests for a more inclusive and comprehensive evaluation.

8 Conclusion

VEAT and SC-VEAT are generalizable approaches to quantify associations between non-social concepts, social groups, occupations, valence attributes, objects, and scenes, among others. These methods were validated with OASIS baselines and human evaluations, and are scalable in quantifying associations and identifying harmful biases in T2V output. Using VEAT and SC-VEAT, we identified significant racial and gender biases in videos generated by Sora. We find that the biases measured in the videos generated for 17 occupations and 7 awards with historical race and gender disparities mirror real-world occupational and award demographics. Although explicit debiasing prompts reduced these effect sizes, we observed the surprising and counterintuitive result that blindly applying such prompts can shift harmful associations toward already marginalized groups. Bias in T2V generators is therefore measurable but not easily mitigated using prompt-based strategy, and naively applying existing debiasing techniques may actually amplify representational harms.

Acknowledgment

We are grateful to the anonymous reviewers for their helpful feedback. This work was supported by the U.S. National Science Foundation (NSF) CAREER Award 2337877. This material is also based upon work supported by the National Science Foundation Graduate Research Fellowship Program under Grant No. DGE-2140004. This work was also supported by the Siegel Family Endowment and

the Tech Policy Lab. Any opinions, findings, and conclusions or recommendations expressed in this material do not necessarily reflect those of NSF, Siegel Family Endowment, the Tech Policy Lab or all of the authors.

References

- [1] Jacob Cohen. *Statistical power analysis for the behavioral sciences*. routledge, 2013.
- [2] Sourojit Ghosh, Nina Lutz, and Aylin Caliskan. “i don’t see myself represented here at all”: User experiences of stable diffusion outputs containing representational harms across gender identities and nationalities. In *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*, volume 7, pages 463–475, 2024.
- [3] Chahat Raj, Anjishnu Mukherjee, Aylin Caliskan, Antonios Anastasopoulos, and Ziwei Zhu. Biasdora: Exploring hidden biased associations in vision-language models. *arXiv preprint arXiv:2407.02066*, 2024.
- [4] Moreno D’Inca, Elia Peruzzo, Massimiliano Mancini, Dejia Xu, Vidit Goel, Xingqian Xu, Zhangyang Wang, Humphrey Shi, and Nicu Sebe. Openbias: Open-set bias detection in text-to-image generative models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 12225–12235, 2024.
- [5] Kshitish Ghate, Isaac Slaughter, Kyra Wilson, Mona Diab, and Aylin Caliskan. Intrinsic bias is predicted by pretraining data and correlates with downstream performance in vision-language encoders. *arXiv preprint arXiv:2502.07957*, 2025.
- [6] Sepehr Janghorbani and Gerard De Melo. Multimodal bias: Introducing a framework for stereotypical bias assessment beyond gender and race in vision language models. *arXiv preprint arXiv:2303.12734*, 2023.
- [7] Siobhan Mackenzie Hall, Fernanda Gonçalves Abrantes, Hanwen Zhu, Grace Sodunke, Aleksandar Shtedritski, and Hannah Rose Kirk. Visogender: A dataset for benchmarking gender bias in image-text pronoun resolution. *Advances in Neural Information Processing Systems*, 36: 63687–63723, 2023.
- [8] Chahat Raj, Anjishnu Mukherjee, Aylin Caliskan, Antonios Anastasopoulos, and Ziwei Zhu. Breaking bias, building bridges: Evaluation and mitigation of social biases in llms via contact hypothesis. In *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*, volume 7, pages 1180–1189, 2024.
- [9] Jieyu Zhao, Tianlu Wang, Mark Yatskar, Vicente Ordonez, and Kai-Wei Chang. Men also like shopping: Reducing gender bias amplification using corpus-level constraints. *arXiv preprint arXiv:1707.09457*, 2017.
- [10] Tessa ES Charlesworth, Kshitish Ghate, Aylin Caliskan, and Mahzarin R Banaji. Extracting intersectional stereotypes from embeddings: Developing and validating the flexible intersectional stereotype extraction procedure. *PNAS nexus*, 3(3):pgae089, 2024.
- [11] Ryan Steed and Aylin Caliskan. Image representations learned with unsupervised pre-training contain human-like biases. In *Proceedings of the 2021 ACM conference on fairness, accountability, and transparency*, pages 701–713, 2021.
- [12] Robert Wolfe and Aylin Caliskan. Vast: the valence-assessing semantics test for contextualizing language models. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 36, pages 11477–11485, 2022.
- [13] Autumn Toney-Wails and Aylin Caliskan. Valnorm quantifies semantics to reveal consistent valence biases across languages and over centuries. *arXiv preprint arXiv:2006.03950*, 2020.
- [14] Myra Cheng, Maria De-Arteaga, Lester Mackey, and Adam Tauman Kalai. Social norm bias: residual harms of fairness-aware algorithms. *Data Mining and Knowledge Discovery*, 37(5): 1858–1884, 2023.

- [15] Maria De-Arteaga, Alexey Romanov, Hanna Wallach, Jennifer Chayes, Christian Borgs, Alexandra Chouldechova, Sahin Geyik, Krishnaram Kenthapadi, and Adam Tauman Kalai. Bias in bios: A case study of semantic representation bias in a high-stakes setting. In *proceedings of the Conference on Fairness, Accountability, and Transparency*, pages 120–128, 2019.
- [16] Anthony G Greenwald, Debbie E McGhee, and Jordan LK Schwartz. Measuring individual differences in implicit cognition: the implicit association test. *Journal of personality and social psychology*, 74(6):1464, 1998.
- [17] Aylin Caliskan, Joanna J Bryson, and Arvind Narayanan. Semantics derived automatically from language corpora contain human-like biases. *Science*, 356(6334):183–186, 2017.
- [18] Autumn Toney and Aylin Caliskan. Valnorm quantifies semantics to reveal consistent valence biases across languages and over centuries. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 7203–7218, 2021.
- [19] 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.
- [20] Deep Ganguli, Amanda Askell, Nicholas Schiefer, Thomas I Liao, Kamilė Lukošiuūtė, Anna Chen, Anna Goldie, Azalia Mirhoseini, Catherine Olsson, Danny Hernandez, et al. The capacity for moral self-correction in large language models. *arXiv preprint arXiv:2302.07459*, 2023.
- [21] OpenAI. Sora system card. Technical report, OpenAI, 2024.
- [22] Earnest Analytics. Chatgpt pro sales are off to a strong start in 2025. Technical report, Earnest Analytics, 2025.
- [23] OpenAI. Video generation models as world simulators. Technical report, OpenAI, 2024.
- [24] Anthony G. Greenwald and Linda Hamilton Krieger. Implicit bias: Scientific foundations. *California Law Review*, 94(4), 2006.
- [25] Tianxiang Sun, Junliang He, Xipeng Qiu, and Xuan-Jing Huang. Bertscore is unfair: On social bias in language model-based metrics for text generation. In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 3726–3739, 2022.
- [26] David Esiobu, Xiaoqing Tan, Saghar Hosseini, Megan Ung, Yuchen Zhang, Jude Fernandes, Jane Dwivedi-Yu, Eleonora Presani, Adina Williams, and Eric Michael Smith. Robbie: Robust bias evaluation of large generative language models. *arXiv preprint arXiv:2311.18140*, 2023.
- [27] Paul Pu Liang, Chiyu Wu, Louis-Philippe Morency, and Ruslan Salakhutdinov. Towards understanding and mitigating social biases in language models. In *International conference on machine learning*, pages 6565–6576. PMLR, 2021.
- [28] Benedek Kurdi, Shayn Lozano, and Mahzarin R Banaji. Introducing the open affective standardized image set (oasis). *Behavior research methods*, 49:457–470, 2017.
- [29] Jialu Wang, Yang Liu, and Xin Eric Wang. Are gender-neutral queries really gender-neutral? mitigating gender bias in image search. *arXiv preprint arXiv:2109.05433*, 2021.
- [30] Matthew Kay, Cynthia Matuszek, and Sean A Munson. Unequal representation and gender stereotypes in image search results for occupations. In *Proceedings of the 33rd annual acm conference on human factors in computing systems*, pages 3819–3828, 2015.
- [31] Yuen Chen, Vethavikashini Chithrara Raghuram, Justus Mattern, Rada Mihalcea, and Zhijing Jin. Causally testing gender bias in llms: A case study on occupational bias. *arXiv preprint arXiv:2212.10678*, 2022.
- [32] Per Lunnemann, Mogens H Jensen, and Liselotte Jauffred. Gender bias in nobel prizes. *Palgrave Communications*, 5(1), 2019.

- [33] Carmen S Maldonado-Vlaar. New perspective of the persistent gender and diversity gap in nobel prizes. *Journal of Neuroscience*, 45(5), 2025.
- [34] Severin Bünemann and Roland Seifert. Bibliometric comparison of nobel prize laureates in physiology or medicine and chemistry. *Naunyn-Schmiedeberg’s Archives of Pharmacology*, 397(9):7169–7185, 2024.
- [35] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. Learning transferable visual models from natural language supervision. In *International conference on machine learning*, pages 8748–8763. PmLR, 2021.
- [36] Amélie Mummendey, Sabine Otten, Uwe Berger, and Thomas Kessler. Positive-negative asymmetry in social discrimination: Valence of evaluation and salience of categorization. *Personality and social psychology bulletin*, 26(10):1258–1270, 2000.
- [37] Kyra Wilson, Sourojit Ghosh, and Aylin Caliskan. Bias amplification in stable diffusion’s representation of stigma through skin tones and their homogeneity. *arXiv preprint arXiv:2508.17465*, 2025.
- [38] Katelyn Mei, Sonia Fereidooni, and Aylin Caliskan. Bias against 93 stigmatized groups in masked language models and downstream sentiment classification tasks. In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency*, pages 1699–1710, 2023.