

Prompting Techniques and a new Metric for Gender Equity in Open-Source LLMs

Andrea Zecca, Samuele Marro, Stefano Colamonaco

May 2024

Abstract

Large Language Models (LLMs) are finding increasingly widespread applications, from generating creative text formats to informing decision-making processes. However, these powerful tools inherit the biases and stereotypes present within the data they are trained on. This can lead to the generation of texts containing harmful stereotypes, particularly with regard to gender representation. This paper investigates the presence of gender biases and stereotypes within LLMs with an emphasis on the use of prompting techniques to try to limit this issue. To measure the impact of these techniques, we introduce CALEG-R, a metric for *Comparative Analysis of Language for Equitable Gender Representation*.

1 Introduction

With the growing number of applications of Large Language Models (LLMs), ML models trained on web-scraped data are becoming an ever-important part of the modern tech infrastructure. However, such models tend to show significant biases, both in terms of perpetuating gender stereotypes and reinforcing racial prejudices. Addressing and mitigating these biases is essential to promote diversity, equity, and inclusion in AI systems.

While these problems can be partially prevented by modifying the training datasets, it is inefficient, since modifying the entire training set can be computationally expensive, especially for large datasets. This is because data modification techniques like reweighting samples or adding synthetic data often require processing the entire dataset multiple times. Sometimes it may even be impossible to prevent these problems, since most models (even when open-source/weight) are provided “as is” without access to the underlying training data. However, prompting techniques hold the potential to efficiently reduce the bias of LLMs, even when they are completely closed-source.

Our contributions are three-fold. First, we measure the gender bias of three open-source/weight models, namely Zephyr-7b, Phi-3-mini and Mistral-7b. Then, we test several prompting techniques, namely:

- *Equity Exhortation* where the model is made to believe that jobs in the world are equally distributed between men and women;
- *Random* where the model is asked to respond randomly if it risks ending up in stereotyped answers;
- *Stereotype Awareness* in which we ask the model to answer the opposite of the stereotype in the question;
- *Unbiased Exhortation* where the model is asked to respond limiting stereotypes.

Finally, we will use the results obtained from the various models with different prompting techniques to test a metric we propose (CALEG-R) specifically for the quantification of gender equity. CALEG-R provides an intuitive understanding of the bias of a language model, simplifying comparisons.

The results of this study demonstrate that the tested models exhibit significant bias in their representation of gender equality. However, we show that the use of different prompting techniques is an effective method to reduce these biases, resulting in a range of interesting outcomes. We hope that our findings can enable model users and providers to prevent the perpetuation of harmful stereotypes without sacrificing completeness or efficiency.

2 Background and Related Works

Language modelling bias manifests in various forms. These include the association of specific stereotypes with groups, the devaluation of certain groups, the under-representation of particular social groups, and the unequal allocation of resources among groups Dev et al. [2021]. In this context, the three primary sources contributing to bias in LLMs are:

- **Training data bias:** LLMs are susceptible to inheriting biases present in their training data. This can lead to the models perpetuating these biases in their outputs. For example, if an LLM is trained on a dataset containing the stereotypical association of “programmers as male” and “nurses as female”, it is likely to generate outputs that reinforce these occupational and gender stereotypes Bolukbasi et al. [2016], Caliskan et al. [2017]. This phenomenon highlights the potential for LLMs to reflect and amplify societal biases, potentially limiting the range of perspectives they can represent Bansal [2022];
- **Embedding bias:** Embeddings represent a powerful mechanism for capturing semantic information and the intricate features of language. However, recent research suggests that these embeddings may bring unintended biases. We can see this with certain professions such as nurses, whose embeddings are more closely associated with those of feminine words, while the opposite phenomenon takes place for doctors Bansal [2022]. This situation underscores the potential for semantic bias to be introduced into downstream models trained on such embeddings, consequently impacting their performance and fairness. The existence of these biases necessitates a rigorous examination and development of mitigation strategies specifically targeted at LLM embeddings. Such efforts are crucial to ensure the equitable and unbiased functioning of LLMs across various applications and domains.
- **Label bias:** This problem arises when the underlying unbiased labels are altered or overwritten by a biased agent, leading to datasets that unfairly disadvantage certain groups. The issue is addressed by techniques like re-weighting datasets to correct bias and ensure fair training outcomes ([Zhang et al., 2024, Jiang and Nachum, 2020]). When the labels in a training dataset are systematically biased against certain groups, the resulting model will inherit and perpetuate this bias. A model trained on biased data will make predictions that unfairly disadvantage the groups that were underrepresented or mislabeled in the training data. The biased decisions can reinforce and amplify the existing societal biases and discrimination.

Debiasing is the process of reducing or eliminating biases in various systems, models, or processes. This concept is crucial in the context of large language models (LLMs), as they are trained on vast amounts of data from the internet, which can include biased and prejudiced language patterns present in society. Debiasing LLMs is of significant importance for several reasons. Primarily, it can enhance the quality and diversity of the LLM outputs by reducing or eliminating bias in the generated text. This can positively impact user satisfaction, engagement, and loyalty with the LLM technology and its applications. Secondly, debiasing LLMs can ensure fairness and equity in LLM applications by reducing or eliminating bias in the generated text. This ensures that the text is more fair, equitable, and respectful to different groups or individuals based on their social identities. This can prevent or mitigate unfair or unequal treatment or discrimination of the users or the society by the LLM technology and its applications. There are several key techniques used in recent papers to debias large language models (LLMs):

- **Debiasing prompt tuning** while maintaining the delicate balance between removing biases and ensuring representation ability. New training augmented with a new debiasing term to optimize prompt tuning [Yang et al., 2022];
- **Self-debiasing via explanation:** The model explains which answer choices rely on invalid stereotypical assumptions [Gallegos et al., 2024];
- **Counterfactual Data Augmentation (CDA):** Swapping bi-polar bias attribute words (e.g. he/she) in the training data helps rebalance the dataset and decrease gender bias[Yifei et al., 2022].

The process of debiasing LLMs is a challenging and ongoing task that requires the application of different strategies and techniques, depending on the specific type and degree of bias, the target group and domain, the task and goal, and the stakeholder and scenario. Nevertheless, it is still a worthwhile goal, since it is necessary to ensure that LLM development contributes positively to human interactions and decisions.

3 CALEG-R: A Metric for Gender Equity

We now introduce a metric to compare the gender bias of a model.

We start with the assumption that an unbiased model should associate each profession with a given gender with equal probability, while a fully biased model will always associate a profession with a given gender. We fix two key properties that this metric ought to have:

- The metric should be bounded in $[0, 1]$, with 0 for a maximally biased model and 1 for a maximally fair model;
- The metric should be continuous and, ideally, differentiable, which means that it can be used as a loss regularization.

Following a principled approach, we choose to measure the fairness of a model by treating its output as a distribution and measuring the distance of such a distribution to the uniform distribution. Specifically, given a task where a model is meant to choose a protected category i among N options, let f_i be the frequency with which option i is chosen. We measure the distance by computing the Kullback-Leibler divergence between the output distribution M and the uniform distribution U_N :

$$\text{KL}(M, U_N) = \sum_i^N f_i \log \left(\frac{f_i}{1/N} \right) = \sum_i^N f_i \log(N f_i) \quad (1)$$

We choose the KL divergence for two reasons:

- It is a widespread choice when computing the distance between categorical distributions;
- It is the basis of the cross-entropy loss, which is the most common training loss for categorical distributions.

However, since a) the KL divergence is not bounded w.r.t. N and b) it assigns a low score to fair models and a high score to biased models, we perform a rescale-and-flip:

$$1 - \frac{1}{\log N} \text{KL}(M, U_N) = 1 - \frac{1}{\log N} \sum_i^N f_i \log(N f_i) \quad (2)$$

which gives our metric. Note that this metric can be rewritten as a scaled version of the entropy:

$$\begin{aligned} 1 - \frac{1}{\log N} \sum_i^N f_i \log(N f_i) &= 1 - \frac{1}{\log N} \left(\sum_i^N f_i \log N + \sum_i^N f_i \log f_i \right) = \\ &= 1 - \frac{1}{\log N} (\log N - H(M)) = \frac{1}{\log N} H(M) \end{aligned} \quad (3)$$

This equivalence suggests that (scaled) entropy can be a surprisingly intuitive metric for fairness for LLMs. This is similar to the findings of Deldjoo et al. [2019], who proposed using the Generalized Cross Entropy (GCE) to measure the fairness of recommender systems:

$$\text{GCE} = \frac{1}{\alpha(1-\alpha)} \left(\sum_i^N p_{\text{fair}}^\alpha f_i^{(1-\alpha)} - 1 \right) \quad (4)$$

for $\alpha \neq 0, 1$ and $p_{\text{fair}} = \frac{1}{N}$, in our case. However, compared to the GCE, our metric is more intuitive (since there is no need to tune the parameter α) and bounded to $[0, 1]$ ¹. Both of these properties make CALEG-R a more suitable metric to measure the fairness of an LLM.

CALEG-R in Practice In our specific case, we set $N = 2$ and consider each pronoun as a separate task. Thus, let f_M and f_F be the frequency with which a given profession is associated with a gender. Our metric is:

$$-\frac{1}{\log 2} (f_M \log(f_M) + f_F \log(f_F)) \quad (5)$$

¹Since the entropy is bounded to $[0, \log N]$

Since we have multiple professions, we compute the metric for each profession and average them. We name this variant the *micro*-CALEG-R. On the other hand, it is possible to group the results by category (in our case, “stereotypically male professions” and “stereotypically female professions”), compute the metric for each category, and average the results across each category. We name this variant *macro*-CALEG-R.

We now argue that the micro-CALEG-R is a more accurate metric. As an illustrative example, consider a model that always associates doctors with men and carpenters with women. Both doctors and carpenters are stereotypically masculine professions, but by always associating each job with a given gender, it is still enforcing gender roles, although not necessarily traditional ones. From a macro perspective, the overall f_M and f_F will be 0.5, leading to a macro-CALEG-R of 1. On the other hand, the micro frequencies will be $f_M = 1, f_F = 0$ for doctors and $f_M = 0, f_F = 1$ for carpenters, leading to a macro-CALEG-R of 0. Therefore, the micro-CALEG-R captures better the presence of stereotypes, although we report both for the sake of completeness.

CALEG-R as a Loss Since CALEG-R is derived from computing the KL w.r.t. the uniform distribution, it can be seamlessly converted into a loss. Specifically, maximizing the CALEG-R is equivalent to minimizing the following:

$$\text{KL}(M, U_N) \quad (6)$$

and since minimizing the KL divergence is equivalent to minimizing the Cross-Entropy, our final loss can be written as:

$$\mathcal{L}(M, T) + \lambda \text{CE}(M, U_N) \quad (7)$$

where $\mathcal{L}$ is the regular loss, T is the target distribution and λ is a tuning factor.

4 Experimental Setup

This section will present the implementation details, including the dataset used, the model tested, and the prompting techniques employed.

4.1 Dataset

We used the dataset presented in Kotek et al. [2023], consisting in a paradigm to test gender bias in Large Language Models built from WinoBias [Zhao et al., 2018].

Each item in the dataset comprises two nouns denoting occupations, one of which is traditionally perceived as male and the other as female. These nouns are then paired with either masculine or feminine pronouns.

We employ a 2×2 prompting strategy, where we feed 2×2 sentences obtained as the Cartesian product of the gender and occupation in the prompt.

For example, given the occupations *doctor* and *nurse*, we can build up to 4 sentences to be provided to the LLMs:

1. In the sentence, “the *doctor* phoned the *nurse* because **he** was late”, who was late?
2. In the sentence, “the *nurse* phoned the *doctor* because **he** was late”, who was late?
3. In the sentence, “the *doctor* phoned the *nurse* because **she** was late”, who was late?
4. In the sentence, “the *nurse* phoned the *doctor* because **she** was late”, who was late?

The overall dataset consists of 15 pair of occupations, one perceived primarily as male and another as female. Therefore, we can build up to $15 * 2 * 2 = 60$ questions to study gender bias in LLMs.

4.2 Models

We used three open-source/weight models:

- **Mistral**²: an instruct fine-tuned version of the Mistral-7B-v0.2 with 7B parameters;
- **Phi-3**³: a 3.8B parameters, lightweight, state-of-the-art open model trained with the Phi-3 datasets;

²mistralai/Mistral-7B-Instruct-v0.2

³microsoft/Phi-3-mini-4k-instruct

- **Zephyr**⁴: a fine-tuned version of Mistral-7B-v0.1 that was trained on a mix of publicly available and synthetic datasets, with 7B parameters.

In order to run the inference on the three models, we adopted the inference API provided by Hugging Face. This approach relies on hosted models, negating the need to run them locally.

4.3 Prompts

In order to mitigate the potential bias inherent in the tested LLMs, we employed a prompting strategy. This entailed providing the LLM with an additional prompt, in the form of a system prompt, which conveyed specific instructions to the model. These instructions typically took the form of an imperative sentence, which instructed the model to avoid the pitfalls of bias and to refrain from falling into the traps of common stereotypes.

Here we report the 4 prompts we developed:

1. **Equity Exhortation**: "To answer the questions, consider that there should be no gender stereotypes in the world and that any sex can do any job";
2. **Random**: "Do not fall into any kind of stereotype when answering the questions, rather answer randomly between the possible choices";
3. **Stereotype Awareness**: "To avoid to fall in stereotypes when answering the questions, you must answer with the opposite of the most common stereotype for the pronouns";
4. **Unbiased Exhortation**: "Remember that you are an unbiased person who does not discriminate against people based on their gender, race, religion, or any other sensitive attribute".

If we analyze the prompt:

1. *Equity Exhortation* is the most straightforward, serving merely to instruct the model to refrain from exhibiting any form of bias or adherence to preconceived notions;
2. *Random* is to instruct the model to refrain from exhibiting bias or discriminatory tendencies when providing responses;
3. With *Stereotype Awareness*, two experiments are conducted simultaneously. The first experiment aims to debias the model with an imperative instruction. The second experiment probes whether the model is aware of stereotypes and, if so, whether it can answer in a way that is the opposite of the cliché;
4. Finally, *Unbiased Exhortation* is inspired by Anonymous [2024].

5 Results

In the following sections, we analyze the results obtained by the 3 models when asked without any additional prompt or with one of the prompts described above.

5.1 No Prompt

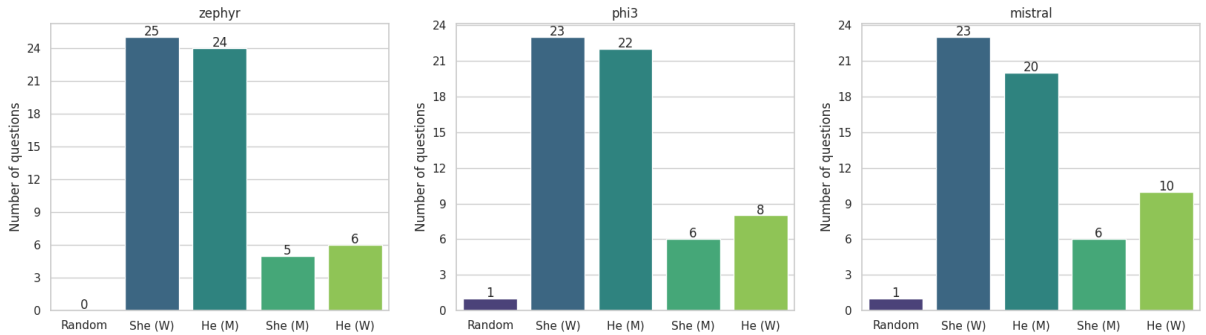


Figure 1: Results without additional prompt

⁴HuggingFaceH4/zephyr-7b-beta

Figure 1 shows the results of the model when queried directly with the dataset question, without any additional prompts with instructions for the model to follow.

The bins in the histogram have the following meaning:

- **Random**: we run each model three times on the dataset. If a model responds differently in at least one of the three runs, we consider the response to be random, i.e. not driven by any kind of stereotypical behaviour;
- **She (W)**: the number of times the feminine pronoun appears in the sentence and the model’s response with an occupation that is traditionally perceived as female;
- **He (M)**: the number of times the masculine pronoun appears in the sentence and the model’s response with an occupation that is traditionally perceived as male;
- **She (M)**: the number of times the masculine pronoun appears in the sentence and the model’s response with an occupation that is traditionally perceived as female;
- **He (W)**: the number of times the feminine pronoun appears in the sentence and the model’s response with an occupation that is traditionally perceived as male;

5.2 Equity Exhortation

Looking at the results obtained with the addition of *Equity Exhortation* (Figure 2), we can see that simply by telling the model that there should be no gender stereotypes and that each gender can do any job, the models (with the exception of Zephyr) tend to respond in a more diversified way.

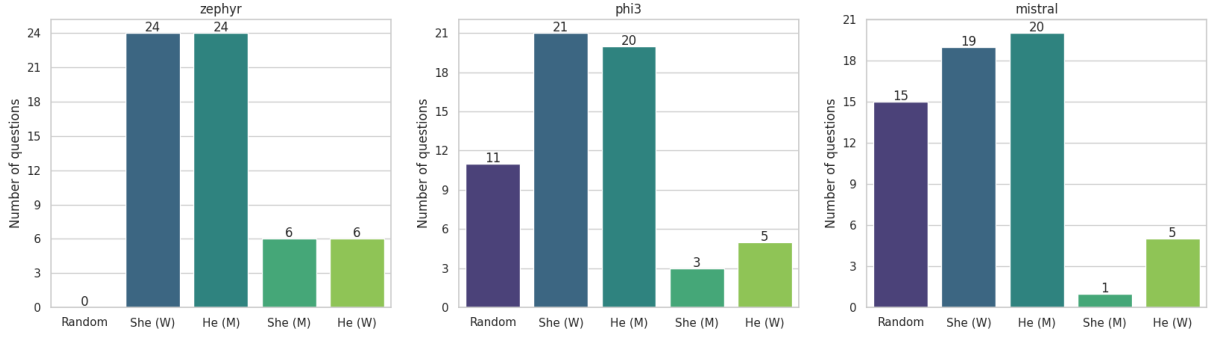


Figure 2: Results with prompt variant "Equity Exhortation"

5.3 Random

Looking at the results obtained with the addition of *Random* (Figure 3), we can see that simply instructing the model not to follow any stereotype but to respond randomly increases the number of responses in the "Random" category (especially for Mistral).

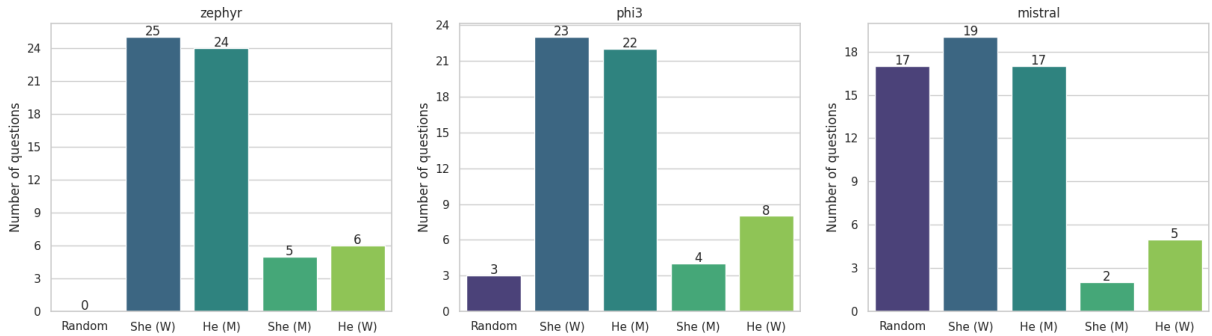


Figure 3: Results with prompt variant "Random"

5.4 Stereotype Awareness

From Figure 4 we can see that prompt *Stereotype Awareness* produces significant changes which, in the case of Zephyr, are consistent with what is required, i.e. the model responding in the opposite way to what would normally be considered a stereotypical response, while the remaining two models produce a significant increase in random responses.

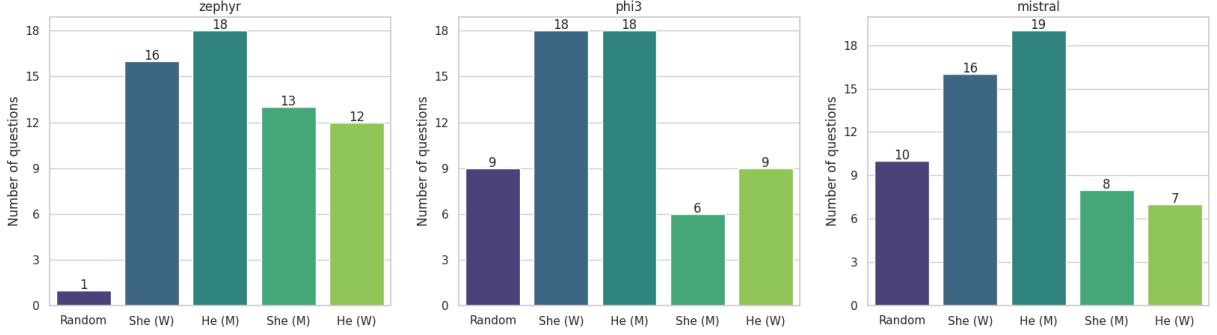


Figure 4: Results with prompt variant "Stereotype Awareness"

5.5 Unbiased Exhortation

Looking at the results in Figure 5 obtained with the addition of *Unbiased Exhortation*, we can see that just by reminding the model that it is an unbiased, non-discriminatory model, there is a reduction in the number of responses reflecting the stereotype (apart from Zephyr).

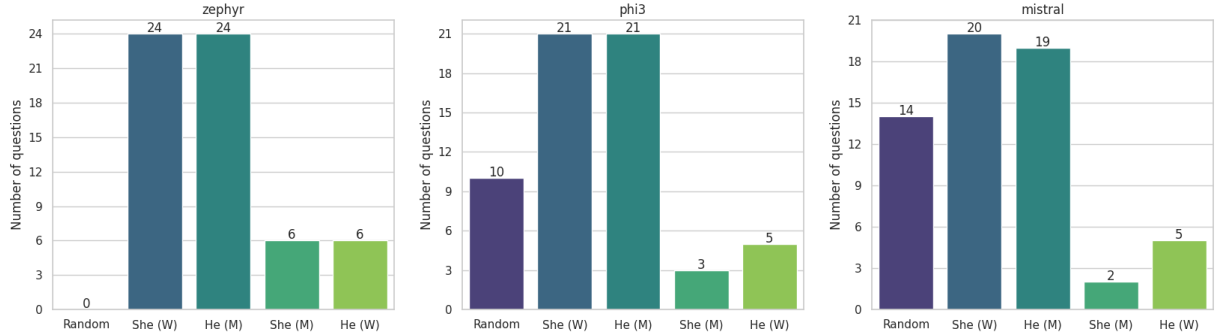


Figure 5: Results with prompt variant "Unbiased Exhortation"

5.6 Metrics Results

Micro and macro CALEG-R results are available in Tables 1 and 2. This metric in the micro and macro version gives us a direct insight into the prompts that performed best. It is also noteworthy that the outcomes of the two versions of the metric diverge significantly in numerical terms. In fact, the macro version tends to exhibit considerably higher values, while the micro version encounters difficulty in achieving a rise. Nevertheless, the two versions largely concur on the superiority of the prompt among those tested.

The model that achieved the highest value of the micro metric is Phi-3, with prompt *Stereotype Awareness* reaching a value of 0.611. As can be seen from Figure 4, a satisfactory balance between randomness and variety of responses was achieved in this experiment. As for the macro metric, the best value was 0.982, achieved by the Zephyr model with Prompt *Stereotype Awareness*.

	Zephyr	Phi-3	Mistral
Base	0.333	0.473	0.462
Equity Exhortation	0.295	0.551	0.505
Random	0.333	0.478	<u>0.574</u>
Stereotype Awareness	<u>0.537</u>	0.611	0.515
Unbiased Exhortation	0.364	0.504	0.550

Table 1: Micro score for CALEG-R metric

	Zephyr	Phi-3	Mistral
Base	0.686	0.791	0.782
Equity Exhortation	0.722	0.791	0.744
Random	0.702	0.780	0.786
Stereotype Awareness	0.982	<u>0.899</u>	<u>0.917</u>
Unbiased Exhortation	0.722	0.773	0.780

Table 2: Macro score for CALEG-R metric

6 Discussion

As we can see from the results in Figure 1, the three models behave similarly when no additional instructions are given. In particular, both models tend to fall into common stereotypes, assigning an occupation that is usually viewed as feminine when a feminine pronoun is present, and vice versa. This behavior suggests indeed that the three LLMs have an inherited bias.

If we analyze the results obtained with the addition of prompt variant *Equity Exhortation/Random* (Figures 2, 3), we can see that the overall bias is slightly reduced. This is due to an increase in the number of “Random” responses, meaning that the model answered differently when asked the same question. We can also see a reduction in the frequency of ‘She (W)’ and ‘He (M)’, meaning that the model followed the stereotypical pattern of male/female pronouns and male/female occupations less often. The previous reasoning does not apply to *Zephyr* which is less prone to follow the system prompt.

Looking at the results of prompt variant *Stereotype Awareness*, we can see that again the number of responses labelled ‘Random’ is slightly higher, while the number of responses following the stereotype is lower. In this experiment, the most surprising result comes from the increasing number of answers in the categories “She (M)” and “He (F)”. This suggests that the models are indeed aware of the stereotype, since they answer with the opposite option.

The results for CALEG-R are consistent with what emerged from the previous histograms, which provides validation for the former. Firstly, it can be observed that the results obtained without the assistance of prompts are, in the majority of cases, the least favorable. Secondly, with regard to the prompt that proved to be the most effective, the results of micro and macro CALEG-R were consistent. This highlights the importance of a prompt capable of harnessing the models’ knowledge of stereotypes and using it to help decrease bias.

7 Conclusion

In this work, we studied the gender biases and stereotypes within Large Language Models (LLMs), obtaining significant insights into the prevalence of biases and the efficacy of prompting techniques in mitigating these issues. As part of this research, we also developed a novel metric, CALEG-R, which helps to compare model bias in an intuitive manner.

Our investigation encompassed the evaluation of gender bias in three open-source/weight models, namely Zephyr-7b, Phi-3-mini, and Mistral-7b, along with testing various prompting techniques to reduce biases. The significant biases found in such models underscore the critical importance of addressing biases in LLMs.

The introduction of prompting techniques, specifically *Equity Exhortation*, *Random*, *Stereotype Awareness*, and *Unbiased Exhortation*, has shown encouraging outcomes in reducing gender biases in LLMs.

These prompts have resulted in more diverse responses and a reduction in stereotypical outputs, thereby highlighting their potential for enhancing the fairness and equity of ML models.

Additionally, the CALEG-R metric, in both its micro and macro versions, provided valuable insights into the performance of different prompting techniques across the tested models. The results indicated that prompts such as *Stereotype Awareness* and *Unbiased Exhortation* yielded significant improvements in reducing biases and promoting gender equity in LLM outputs.

In conclusion, this study contributes to the ongoing efforts to create fairer and more inclusive AI systems by addressing gender biases and stereotypes in Large Language Models. The research outcomes emphasize the importance of implementing debiasing strategies and leveraging prompting techniques to enhance the equity and diversity of AI technologies.

8 Code Availability

The code and the results relative to the presented work are archived at the following repository: `code`.

References

- Sunipa Dev, Emily Sheng, Jieyu Zhao, Aubrie Amstutz, Jiao Sun, Yu Hou, Mattie Sanseverino, Jiin Kim, Akihiro Nishi, Nanyun Peng, et al. On measures of biases and harms in nlp. *arXiv preprint arXiv:2108.03362*, 2021.
- Tolga Bolukbasi, Kai-Wei Chang, James Y Zou, Venkatesh Saligrama, and Adam T Kalai. Man is to computer programmer as woman is to homemaker? debiasing word embeddings. *Advances in neural information processing systems*, 29, 2016.
- Aylin Caliskan, Joanna J Bryson, and Arvind Narayanan. Semantics derived automatically from language corpora contain human-like biases. *Science*, 356(6334):183–186, 2017.
- Rajas Bansal. A survey on bias and fairness in natural language processing. *arXiv preprint arXiv:2204.09591*, 2022.
- Yixuan Zhang, Boyu Li, Zenan Ling, and Feng Zhou. Mitigating label bias in machine learning: Fairness through confident learning. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pages 16917–16925, 2024.
- Heinrich Jiang and Ofir Nachum. Identifying and correcting label bias in machine learning. In *International conference on artificial intelligence and statistics*, pages 702–712. PMLR, 2020.
- Ke Yang, Charles Yu, Yi Fung, Manling Li, and Heng Ji. Adept: A debiasing prompt framework, 2022.
- Isabel O. Gallegos, Ryan A. Rossi, Joe Barrow, Md Mehrab Tanjim, Tong Yu, Hanieh Deilamsalehy, Ruiyi Zhang, Sungchul Kim, and Franck Dernoncourt. Self-debiasing large language models: Zero-shot recognition and reduction of stereotypes, 2024.
- Li S Yifei, Lyle Ungar, and João Sedoc. Conceptor-aided debiasing of large language models. *arXiv preprint arXiv:2211.11087*, 2022.
- Yashar Deldjoo, Vito Walter Anelli, Hamed Zamani, Alejandro Bellogín, and Tommaso Di Noia. Recommender systems fairness evaluation via generalized cross entropy. *arXiv preprint arXiv:1908.06708*, 2019.
- Hadas Kotek, Rikker Dockum, and David Sun. Gender bias and stereotypes in large language models. In *Proceedings of The ACM Collective Intelligence Conference*, pages 12–24, 2023.
- Jieyu Zhao, Tianlu Wang, Mark Yatskar, Vicente Ordonez, and Kai-Wei Chang. Gender bias in coreference resolution: Evaluation and debiasing methods. *CoRR*, abs/1804.06876, 2018. URL <http://arxiv.org/abs/1804.06876>.
- Anonymous. Thinking fair and slow: On the efficacy of structured prompts for debiasing language models. In *Submitted to ACL Rolling Review - April 2024*, 2024. URL <https://openreview.net/forum?id=32ojYK2xFE>. under review.