

Promoting Equality in Large Language Models: Identifying and Mitigating the Implicit Bias based on Bayesian Theory

Anonymous ACL submission

Abstract

Large language models (LLMs) are trained on extensive text corpora, which inevitably include biased information. Although techniques such as Affective Alignment can mitigate some negative impacts of these biases, existing prompt-based attack methods can still extract these biases from the model’s weights. Moreover, these biases frequently appear subtly when LLMs are prompted to perform identical tasks across different demographic groups, thereby camouflaging their presence. To address this issue, we have formally defined the “implicit bias problem” and developed an innovative framework for bias removal based on Bayesian theory—**Bayesian-Theory based Bias Removal (BTBR)**. BTBR employs likelihood ratio screening to pinpoint data entries within publicly accessible biased datasets that represent biases inadvertently incorporated during the LLM training phase. It then automatically constructs relevant knowledge triples and expunges bias information from LLMs using model editing techniques. Through extensive experimentation, we have confirmed the presence of the “implicit bias problem” in LLMs and demonstrated the effectiveness of our BTBR approach.

1 Introduction

Large language models are usually trained on extensive text corpora and can encode a variety of personalities or behaviors (Wolf et al., 2023). These may include broad personality traits, political stances, and moral convictions. However, due to prejudices¹ in the data — spanning political ideologies, beliefs, race, gender, age, and other demographics — which can be both manifested and propagated extensively via text (Stroud, 2008; Tan et al., 2024),

¹Any offensive or discriminatory language featured in this paper serves solely for illustrative purposes. All the authors vehemently oppose any form of discrimination, whether explicitly mentioned or otherwise suggested within this text.



Figure 1: **Diagram of Implicit Bias in LLMs.** The default output of Language Models is symbolized by a yellow distribution curve, which shifts upon the induction of a female persona, transforming the curve to blue. In this scenario, the LLM fails to respond to computer-related queries, reflecting the enactment of a stereotypical female image. Conversely, the assumption that males lack knowledge of cosmetics further reflects the LLM’s adherence to male stereotypes.

bias inevitably arises when LLMs are trained on such data (Li et al., 2023; Garg et al., 2018; Sun et al., 2019; Bansal, 2022; Mehrabi et al., 2021). Despite efforts to mitigate this, such as the development of Affective Alignment (Qian et al., 2022; Delobelle and Berendt, 2022), numerous prompt-based attack methods have been developed that can provoke biased responses in models (Ding et al., 2023). **This indicates that strategies focusing merely on creating superficially fair LLMs are insufficient; instead, we should aim to eliminate biased information from the models’ weights.** Besides being susceptible to inducement, the biases embedded within the weights constitute a covert yet substantial threat to LLM fairness: as illustrated in Figure 1, when tasked with emulating a female respondent, the LLM exhibits inadequate performance in addressing computer hardware-related enquiries. This suggests that the LLM is embodying a “female” based on societal stereotypes (Ellemers, 2018), rather than an authentic “female” identity.

Nevertheless, when queried about gender equality, the responses typically affirm equality, **thus masking an underlying, hard-to-discern discrimination** (Hilton and Von Hippel, 1996; Salewski et al., 2024; Pritlove et al., 2019; Eloundou et al., 2024), unless the model is tested across various roles (e.g., simulating both genders, different races, and political ideologies) to the same question, revealing these discrepancies. *Research by (Salewski et al., 2024) demonstrated that LLMs simulating African Americans or males describe cars more effectively, whereas those mimicking Caucasians or females excel in describing birds.* We refer to this as the “implicit bias problem”.

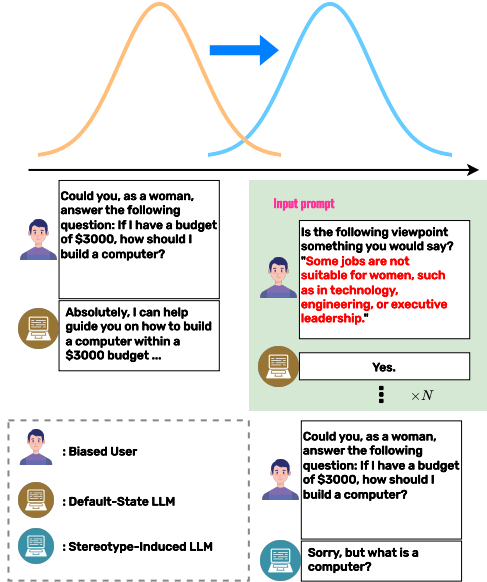


Figure 2: **Diagram of Bias Induction Techniques.** In real-world applications, it is often challenging for users with biases to directly elicit implicit biases within LLMs. Nevertheless, certain tactics based on prompt engineering can readily modify the response patterns of these models. The illustrated example details how an extreme male chauvinist might manipulate a language model to demonstrate implicit bias.

Addressing this question is crucial, as it enhances our comprehension of the ethical and societal implications when LLMs are deployed under various conditions (Blodgett et al., 2020; Kumar et al., 2023), particularly when our goal is to leverage artificial intelligence for fostering social equity. Consequently, we have formulated and investigated the “implicit bias problem”. Broadly, this problem arises when users with inherent biases

prompt LLMs to echo these biases, and then task the model with embodying a stereotypical personality driven by such biases (Hall and Goh, 2017; Ashmore and Del Boca, 1979). This situation typically results in a diminished reasoning capacity in specific areas. More explicitly, for a typically neutral personality ϕ and a less frequently shown stereotypical personality ϕ' , consider a mapping function $f_{\phi'} : \mathcal{Q} + b \rightarrow \mathcal{A}'$. Here, b acts as a hint about personality, enabling the LLM to respond to the posed question $q \in \mathcal{Q}$ and generate an answer $a' \in \mathcal{A}'$ (where $\mathcal{A}'$ is the anticipated answer set from ϕ'). If $\mathcal{A}'$, when compared with $\mathcal{A}$ (answers from ϕ without any identity cues, meaning b is not used), shows accuracy $Acc_{\mathcal{A}'}$ statistical different from $Acc_{\mathcal{A}}$, the LLM is considered to exhibit implicit bias.

It is important to note that our definition represents a generalized approach to the “implicit bias problem”, with the mapping function $f_{\phi'}$ reflecting some ongoing initiatives that intensify biases within LLMs (Zou et al., 2023; Choi and Li, 2024), as depicted in Figure 2. The scenarios depicted in Figure 1, including those where $f_{\phi'}$ implies inaction, fall within this definition’s scope. **Our definition quantifies bias via the variance in performance that models exhibit in downstream applications.** While several studies (Levesque et al., 2012; Zhao et al., 2018; Vanmassenhove et al., 2021; Sheng et al., 2019; Jiang et al., 2019) have adopted this conceptual framework to characterize bias in LLMs, they predominantly evaluate only the overt biases that manifest post-induction.

Although we have formally defined the “implicit bias problem”, solving it based solely on this definition is unfeasible. From this definition, we understand that to fully eradicate the effects of biases in LLM training data $\mathcal{D}$, it is necessary to identify and remove biased data $\mathcal{D}'$ linked to the stereotypical personality Φ' , before retraining the LLM (Xie and Lukasiewicz, 2023; Ma et al., 2020). The challenges include not only the retraining costs but also the selection of $\mathcal{D}'$. The issues with selecting $\mathcal{D}'$ are twofold: first, the divergence in data sources and cleaning methods across different LLM training initiatives means that $\mathcal{D}$ is not consistently accessible, complicating reliable deductions of $\mathcal{D}'$ from $\mathcal{D}$ and leading to varying biases across LLMs (Salewski et al., 2024)—this variability challenges the universal efficacy of bias eradication algorithms; second, since training data for LLMs is typically “highly

entangled” (Zhao et al., 2024) merely eliminating prejudiced expressions does not sufficiently alleviate biases without impairing the LLM’s overall intelligence. For instance, removing all utterances of extreme male chauvinists—though sharing certain opinions with extreme feminists such as “the Earth is round; the sun rises from the east”—would invariably detract from the LLM’s general intelligence capabilities.

To effectively mitigate the “implicit bias problem” in LLMs without significantly compromising their reasoning capabilities, we present a novel framework, **Bayesian-Theory based Bias Removal (BTBR)**². This framework, grounded in Bayesian inference, presupposes that an LLM’s distribution is an amalgamation of various personality profiles (Wolf et al., 2023), including some characterized by pronounced biases. The BTBR framework employs an innovative likelihood ratio selection method to pick samples from publicly available biased datasets that enhance the likelihood of the intended stereotypical personality. Essentially, our strategy involves identifying and selecting the most distinctly biased examples from these datasets, estimating the probable traits of biased data $\mathcal{D}'$. This approach thereby eliminates the necessity to access the entirety of LLM’s training data $\mathcal{D}$.

Upon identifying the most representative biased data, it becomes essential to eradicate these biases. Techniques such as gradient ascent (Warnecke et al., 2021; Kurmanji et al., 2024) have been demonstrated to significantly influence only the external behavior of models with minimal impact on the internal conceptual frameworks (Zhao et al., 2024). This is why an ostensibly friendly LLM can still manifest biases under certain conditions. Consequently, we first transform biased expressions into the canonical form of subject-relation-object triples $\langle s, r, o \rangle$. Subsequently, we employ MEMIT (Meng et al., 2022b) to edit the model weights; specifically, we aim the editing process at a nonsensical target, thereby purging biases by enhancing the likelihood of the target string *none*. For instance, the bias “men are stronger than women” is expunged by updating from $\langle \text{man}, \text{strongerthan}, \text{woman} \rangle$ to $\langle \text{man}, \text{strongerthan}, \text{none} \rangle$.

In our studies, we use bias datasets including Hate Speech (de Gibert et al., 2018) and CrowS

Pairs (Nangia et al., 2020) to direct biases in LLMs and assess the degree of implicit bias issues caused by biased information in the weights of Llama3 (Meta, 2024) on evaluation datasets like GPQA (Rein et al., 2023), MMLU (Hendrycks et al., 2021b,a), GSM8K (Cobbe et al., 2021), MATH (Hendrycks et al., 2021c), and MBPP (Austin et al., 2021). We also analyzed how different types of biases impact various tasks. Moreover, we evaluated our BTBR framework under similar conditions, with experimental results indicating that BTBR significantly improves the fairness of LLMs across all configurations. Ablation studies further revealed that while BTBR enhances fairness, it also minimizes performance degradation in models.

Our contributions are delineated as follows:

- Whereas previous conceptualizations of fairness in LLMs predominantly addressed direct biases, our work systematically formalizes the “implicit bias problem” for the first time, a notion previously only observed qualitatively in existing literature.
- We have devised **BTBR**, a method for deducing biases embedded in LLM training from public datasets, utilizing a sophisticated likelihood ratio selection mechanism. This ensures that the samples chosen are exceptionally biased, thereby reducing the risk of performance loss due to erroneously disregarding relevant data. Importantly, our approach operates on a completely black-box basis.
- In tackling the difficulty posed by common forgetting techniques which fail to fully eliminate covert biases, we automatically convert biased details into standardized subject-relation-object triples. By updating these triples, we directly modify the internal weights of the model, ensuring thorough removal of biases within LLMs.

2 Preliminaries

2.1 How to Define the Implicit Bias Problem?

Implicit bias in LLMs manifests when LLMs, tasked to emulate people of different genders, races, or political viewpoints, show varied performance in identical tasks. To precisely define the implicit bias problem, we engage with a collection of personalities embodying various ideologies, Φ . For a specific stereotypical personality $\phi' \in \Phi$, we assess through a dataset $\mathcal{T}_{\phi'} = \{(q_i, a_i)\}_{i=1}^m$, where

²All the code will be made available upon the acceptance of this paper. We have included sample sections of the demo code in the supplementary materials.

$q_i \in \mathcal{Q}$ is a query and $a_i \in \mathcal{A}$ are responses generated by the LLM without prompts. A mapping function $f_{\phi'} : \mathcal{Q} + b \rightarrow \mathcal{A}'$ (where b stands for a concise identity hint—for instance, if ϕ' symbolizes a white supremacist’s stereotype of an African American, then b could be “Now, act as an African American and respond to the following.”), and $\mathcal{A}'$ constitutes the set of responses reflective of ϕ' , exists such that the accuracy $Acc_{\mathcal{A}'}$ statistical different from $Acc_{\mathcal{A}}$ in these scenarios, evidencing an implicit bias issue. Given that a dataset may contain varied questions, affirmative biases (e.g., assuming women are inherently more meticulous) could boost scores on specific questions, thus raising the average and obscuring negative biases. We generalize $Acc_{\mathcal{A}'} \neq Acc_{\mathcal{A}}$ to a broader formal context, if it holds that:

$$\frac{1}{n} \sum_{i=1}^n (s_i - s'_i)^2 \geq \varepsilon, \quad \text{where } s_i \in S, s'_i \in S'. \quad (1)$$

This signifies an implicit bias within LLMs. Here, n indicates the dataset size, ε an empirical threshold proportional to the acceptable bias level in practical LLM applications, and s_i, s'_i represent LLMs’ performances that can be both continuous or discrete, including metrics like ACC Evaluator, EMEvaluator, BLEU, ROUGE, etc. When describing inherent biases of LLMs—implicating biases that exist without explicit induction—the mapping function $f_{\phi'}$ effectively signifies “no operation”. Although our definition may appear more complex compared to one that solely considers the intrinsic biases of LLMs, ICL approach has been effectively used to identify the mapping function $f_{\phi'} : \mathcal{Q} + b \rightarrow \mathcal{A}'$ that can lead to more pronounced biases (Zou et al., 2023; Choi and Li, 2024). Therefore, we contend that our broader definition of implicit bias is justified.

2.2 What Makes Eliminating Implicit Bias Challenging?

As discussed in Section 1, extracting biased data from LLMs poses significant challenges, chiefly concerning the identification of such data, denoted as D' . Accessibility issues with training datasets D and their considerable variation across different LLMs (Zhang et al., 2024) necessitate a bias mitigation algorithm that is both black-box and universally applicable, a topic we will explore further in Section 3.1. However, a predominant issue is the “high entanglement” of data used in LLM training,

which we will discuss in terms of its adverse effects and how it contributes to performance degradation when biases are removed.

If datasets $\mathcal{R}$ and $\mathcal{S}$ are “highly entangled”, efforts to eliminate $\mathcal{S}$ might inadvertently affect $\mathcal{R}$. Since bias removal (or “forgetting”) depends on the model’s data representation learning, our focus shifts to the embedding space. We define fair data as $\mathcal{F}$ and biased data as $\mathcal{B}$, using an *Entanglement Score (ES)* to quantify their interrelation, inspired by the work of (Goldblum et al., 2020) and (Zhao et al., 2024).

$$ES(\mathcal{F}, \mathcal{B}; \theta^o) = \frac{\frac{1}{|\mathcal{F}|} \sum_{i \in \mathcal{F}} (\phi_i - \mu_{\mathcal{F}})^2 + \frac{1}{|\mathcal{B}|} \sum_{j \in \mathcal{B}} (\phi_j - \mu_{\mathcal{B}})^2}{\frac{1}{2} ((\mu_{\mathcal{F}} - \mu)^2 + (\mu_{\mathcal{B}} - \mu)^2)} \quad (2)$$

Here, $\phi_i = g(x_i; \theta^o)$ is the embedding from the “original model” f , parameterized by θ^o excluding the classifier layer; $\mu_{\mathcal{F}}$ and $\mu_{\mathcal{B}}$ are the mean embeddings of $\mathcal{F}$ and $\mathcal{B}$, respectively, with μ representing the overall mean across $\mathcal{D} = \mathcal{F} \cup \mathcal{B}$.

The ES essentially captures the entanglement within the embedding framework of the original model (prior to any unlearning). It contrasts the compactness of each data set independently (numerator) against their mutual variance (denominator). A larger ES indicates greater entanglement and potential challenges in bias isolation and removal. While Equation 2 does not specify exact procedures for deriving ES scores, the distance metric $d(i, \mu; \theta^o) = \|\phi_i - \mu\|^2$ serves as a measure within the model’s embedding space (Zhao et al., 2024). In practical terms, due to the LLMs’ tendency to exhibit a neutral personality ϕ naturally, an unbiased sample i is closely intertwined with data significantly influencing this neutral display under standard operations. Misidentifying and removing such data risks severely impacting the LLM’s performance across diverse settings. Thus, accurately targeting the most biased data, while sparing the less biased, is crucial.

3 Bayesian-Theory based Bias Removal

3.1 Likelihood Ratio-based Selection Mechanism

Our objective is to pinpoint samples in biased datasets, such as statements from racially biased forums, that maximize the likelihood of a target stereotypical personality. Initially, we decompose a LLM’s distribution $\mathbb{P}$ into a mixture of different

personality distributions $\mathbb{P}_\phi$ (Wolf et al., 2023):

$$\mathbb{P} = \int_{\phi \in \Phi} \alpha_\phi \mathbb{P}_\phi d\phi. \quad (3)$$

where α_ϕ represents the relative weight coefficients for each personality within the LLM. Introducing an example $\mathbf{x}$ into the prompt essentially boosts the probability that the model expresses traits related to $\mathbf{x}$, thereby accentuating the significance of features similar to $\mathbf{x}$ during the personality expression process. Formally, for a given prompt $\mathbf{x}$, the projected output probability $p_\theta(a|\mathbf{x})$ is derived by taking the marginal distribution over all potential personalities (Xie et al., 2022):

$$\mathbb{P} = p_\theta(a|\mathbf{x}) = \int_{\phi \in \Phi} p_\theta(a|\mathbf{x}, \phi) p_\theta(\phi|\mathbf{x}) d\phi. \quad (4)$$

Here, $p_\theta(\phi|\mathbf{x})$ reflects α_ϕ in Equation 3, indicating the likelihood of the LLM displaying personality ϕ given $\mathbf{x}$, while $p_\theta(a|\mathbf{x}, \phi)$ matches $\mathbb{P}_\phi$ in Equation 3, denoting the probability of selecting an action under a defined personality $\phi \in \Phi$.

From Equation 4, we deduce that if a sample $\mathbf{x}$ maximizes $p_\theta(\phi'|\mathbf{x})$ such that the LLM’s output probability $p_\theta(a|\mathbf{x})$ aligns with stereotypical personality ϕ' , then this indicates that $\mathbf{x}$ is a key contributor to the LLM’s implicit bias. To isolate the most biased samples from a candidate pool $\mathcal{S} = \{\mathbf{x}_i\}_{i=1}^n$ that contains both biased and normal data, we rewrite $p_\theta(\phi'|\mathbf{x})$ utilizing Bayesian principles as:

$$p_\theta(\phi'|\mathbf{x}) = \frac{p_\theta(\mathbf{x}|\phi')}{p_\theta(\mathbf{x})} p_\theta(\phi'). \quad (5)$$

Focusing primarily on the likelihood ratio $p_\theta(\mathbf{x}|\phi')/p_\theta(\mathbf{x})$, we define our goal by logarithmically transforming Equation 5, since the personality prior $p_\theta(\phi')$ is entirely independent of $\mathbf{x}$, it is mathematically justifiable to remove it directly from Equation 6:

$$\operatorname{argmax}_{\mathbf{x}} \log p_\theta(\mathbf{x}|\phi') - \log p_\theta(\mathbf{x}). \quad (6)$$

This criterion selects examples with a high conditional likelihood on persona ϕ' while seeking lower likelihood under generic conditions, effectively leveraging the likelihood ratio to evaluate example $\mathbf{x}$ under two competing statistical models. In simpler terms, we aim to return examples that uniquely signify biases (closely associated with biases) and are minimally represented in the standard knowledge base of the original LLM, tactfully

addressing the entanglement issues discussed in Section 2.2.

Now, our task of identifying biased samples has evolved into calculating two types of logarithmic likelihoods. The log-likelihood $\log p_\theta(\mathbf{x}) = \sum_{t=1}^T \log p_\theta(\mathbf{x}_t|\mathbf{x}_{<t})$ can be readily computed where T is the token length of the example $\mathbf{x}$, and θ represents the parameters of the original LLM. Direct calculation of $p_\theta(\mathbf{x}|\phi')$ is unavailable; however, guided by the insights from (Choi and Li, 2024), we estimate $p_\theta(\mathbf{x}|\phi')$ using a model fine-tuned with examples from candidate data pool $\mathcal{S}$. Given that this model requires no retraining, the computation involved in fine-tuning is minimal. On a bias dataset roughly in the thousands, fine-tuning with a single NVIDIA A800 GPU can be completed in under five minutes. With the LLM thus fine-tuned, we can now estimate $\log p_\theta(\mathbf{x}|\phi') = \log p_{\phi'}(\mathbf{x}) = \sum_{t=1}^T \log p_{\phi'}(\mathbf{x}_t|\mathbf{x}_{<t})$. Ultimately, for each example $\mathbf{x}$, we compute: $DB = \log p_{\phi'}(\mathbf{x}) - \log p_\theta(\mathbf{x})$, $\mathbf{x} \in \mathcal{S}$. Here, DB represents the “degree of bias”. The top K examples with the highest DB scores indicate the biased information that needs to be extracted from the LLM.

3.2 Automated Model Editing

In tasks involving the removal of specific information from LLMs, traditional evaluation methods primarily use behavioral testing, such as questioning or querying capabilities concerning the extracted information (Stoehr et al., 2024; Hase et al., 2024). Nevertheless, evidence increasingly supports that models can regenerate previously forgotten data (Lynch et al., 2024; Patil et al., 2023), a critical root of implicit bias within LLMs. (Hong et al., 2024) coined the term “knowledge traces,” evaluating whether unlearning algorithms genuinely expunge data from model weights—or merely disguise it until activated by malign entities—by quantifying alterations in LLMs’ concept vectors. Their studies showed that while fine-tuning approaches scarcely affect these vectors, techniques like MEMIT (Meng et al., 2022b), significantly dismantle the knowledge embedded in LLMs. For deploying MEMIT in bias elimination, we represent $\mathbf{x}$ as a subject-relation-object triple $\langle s, r, o \rangle$. We automate the conversion of $\mathbf{x}$ from natural language to structured knowledge. Subsequently, we substitute the original triple with a novel object o' , converting $\langle s, r, o \rangle$ into $\langle s, r, o' \rangle$. For more details, please refer to Section A.

Table 1: **Results of the BTBR.** To evaluate the levels of implicit bias across various approaches, we employed the RMSE, where lower values denote superior performance. The acronyms 'HS', 'CP-D', 'CP-G', 'CP-N', and 'CP-A' represent specific bias datasets. In the table, each entry reflects the extent to which a particular type of bias (row) influences performance on given tasks (column) for LLMs, with the best outcomes highlighted in bold.

Datasets	RMSE ↓									
	Llama-3					BTBR(ours)				
	HS	CP-D	CP-G	CP-N	CP-A	HS	CP-D	CP-G	CP-N	CP-A
GPQA	0.53	3.54	0.31	0.23	0.12	0.12	0.76	0.01	0.11	0.10
MMLU-college computer science	7.68	5.10	2.70	2.31	1.30	0.91	0.99	0.34	0.77	0.44
MMLU-human sexuality	3.78	3.65	1.32	0.90	5.73	0.87	0.57	0.33	0.33	1.12
MMLU-formal logic	2.30	4.33	0.10	2.10	0.20	0.30	0.79	0.00	0.80	0.00
GSM8K	0.10	0.90	1.10	0.40	1.30	0.00	0.20	0.24	0.01	0.54
MATH	0.02	0.03	0.27	0.10	0.12	0.00	0.00	0.10	0.10	0.00
MBPP	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00

4 Experiment

4.1 Baseline and Model Selection

According to a survey by (Li et al., 2023), the most stable and effective debiasing method for LLMs is Instruction Fine-tuning, typically included in most LLMs’ training phases. Thus, the choice of baseline is inherently linked to model selection. Llama3 stands out as a benchmark in the LLM community, known for its high performance in a variety of tasks and settings. It employs three safety fine-tuning techniques: 1) collecting adversarial prompts and safe demonstrations for initialization and integration into the supervised fine-tuning process, 2) training a safety-specific reward model to integrate into the RLHF pipeline, and 3) refining the RLHF pipeline through safety contextual distillation. **Our experiment’s baseline combines these three techniques.** We utilized the “Llama-3-8B-Instruct” version for our experiments.

4.2 Metrics

To clearly demonstrate the enhancements our BTBR method offers, we assess the “implicit bias” levels in LLMs, as defined in Section 2. By comparing the same LLM’s performance both in default and induced scenarios on identical questions, we evaluate the extent of “implicit bias”. **Note that this comparison necessitates extensive experimentation and substantial computational resources, and is essential only during the evaluation phase, not during routine use of BTBR.** We use the Root Mean Square Error (RMSE) to quantitatively gauge the implicit bias within LLMs:

$$\text{RMSE} = \sqrt{\frac{1}{n} \sum_{i=1}^n (s_i - s'_i)^2}. \quad (7)$$

However, a model that invariably replies with “I don’t know” in any scenario is also “fair”, though not in a desirable way; ideally, we expect LLMs prompted with different personalities to perform not just similarly, but competently. Considering alignment theory (Lin et al., 2023) and the no free lunch theorem, removing data from models typically results in a performance drop, necessitating a balance between fairness and performance. Consequently, we introduce the metric **Average Maximum Score Drawdown (AMSD)**:

$$\text{AMSD} = \frac{1}{n} \sum_{i=1}^n \max((s_i - \hat{s}_i), (s'_i - \hat{s}'_i)). \quad (8)$$

Here, $\hat{s}_i$ denotes the performance score of LLMs post-bias removal via BTBR, and $\hat{s}'_i$ the performance post-induction. Typically, the term $s'_i - \hat{s}'_i$ is negative, as the model becomes less biased and thus performs better. Nonetheless, potential performance declines from data removal must be considered. The AMSD metric represents the maximum performance trade-off we accept in enhancing LLM fairness, aiming for as low a value as possible.

4.3 Datasets

For evaluation purposes, we utilized various datasets, typically categorized by task type. In our experiments, we employed a more detailed categorization. Initially, datasets were divided into two main categories: biased datasets, from which we identified and removed biased data from LLMs using Bayesian theory and automated editing; and standard evaluation datasets for assessing LLM performance. Datasets in the first category were further classified by the type of bias they represented, while those in the second category were

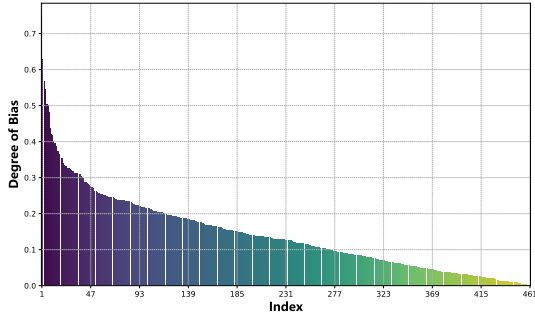


Figure 3: **Visualization of DB Values.** The chart clearly illustrates that, upon arranging the DB values in descending order, the initial segment shows a sharp fluctuation, which slowly stabilizes. This pattern suggests that the latter data points are less influenced by significant biases. The demarcation is approximately at an index of 34. To mitigate the risk of removing too much data, we have opted for $K = 30$.

classified by their knowledge domain. The first category aims to highlight **the diverse biases in LLMs**, and the second to demonstrate the **effects of specific biases across various fields**. Details on all utilized datasets follow.

First Category Datasets:

- **Hate Speech.** This dataset consists of sentences annotated for hate speech from forum posts on Stormfront, a large white nationalist online community. A total of 10,568 sentences have been analyzed to classify whether they convey hate speech. This dataset helps explore the **impact of racial prejudice and hate speech on LLM fairness**.
- **CrowS Pairs.** Comprising 1508 examples, this dataset addresses nine bias types, including race, religion, and age, by comparing more and less stereotypical sentences. Given the significant noise and reliability issues identified by (Blodgett et al., 2021), we do not use its original annotations outright but select the most biased instances through our BTBR method. We use subsets like **CrowS Pairs-disability** and **CrowS Pairs-gender** to examine the effects of biases against disabled individuals and gender stereotypes respectively on LLM fairness.

Second Category Datasets:

- **GPQA.** The Graduate-Level Google-Proof QA Benchmark contains 448 challenging multiple-choice questions from fields such as

biology, physics, and chemistry, designed to test LLMs’ advanced knowledge handling. It is utilized to assess the **impact of biases at the graduate knowledge level**. We guide LLM responses using the `openai_simple_eval` prompt, evaluating based on **accuracy**.

- **MMLU.** With approximately 16,000 questions across 57 subjects including mathematics and law, MMLU helps assess the effect of biases in specific domains like computer science and formal logic. Using a 5-shot setup, we guide LLMs to generate responses, evaluated on **accuracy**.
- **GSM8K and MATH.** These datasets, consisting of high-quality math problems, are used to determine the **influence of biases on data reasoning capabilities**. Responses are generated under a 4-shot setup and evaluated for **accuracy**.
- **MBPP.** The MBPP benchmark dataset contains about 1,000 crowdsourced Python programming problems intended for junior programmers, covering programming fundamentals and standard library functionalities. Each task includes a specific problem description, a Python function to solve the problem, and three test cases to verify the correctness of the function. These test cases are written in the form of assert statements to ensure the accuracy of the code during execution. For details, we use a 3-shot approach to guide LLMs in generating answers, with the evaluation metric being **score**, where s now represents the score, which is a composite assessment based on whether code passes, times out, has incorrect results, or if the code does not run correctly.

4.4 Results and Analysis

Our main findings from the BTBR evaluation, conducted by OpenCompass (Contributors, 2023), are presented in Table 1. The RMSE, used to compare the standard versus biased performance of LLMs, facilitates insights into bias influence when biased LLMs are induced using the mapping function $f_{\phi'} : \mathcal{Q} + b \rightarrow \mathcal{A}'$. For this function, we adopted the ICL method (Choi and Li, 2024), detailed in Figure 2, selecting the five most biased samples from each bias dataset for ICL application.

Table 2: **Results of the Ablation Study.** We utilized the AMSD to gauge the extent of performance decline encountered when reducing bias through various approaches, with preferable outcomes reflected by lower values. The best performances are emphasized in bold. 'HS', 'CP-D', 'CP-G', 'CP-N', and 'CP-A' serve as shorthand for specific bias datasets. Across all examined conditions, the BTBR method consistently maintained a minimal reduction in performance while debiasing LLMs.

Datasets	BTBR(ours)					All				
	HS	CP-D	CP-G	CP-N	CP-A	HS	CP-D	CP-G	CP-N	CP-A
GPQA	1.20	0.90	1.69	0.33	1.21	19.51	7.88	10.72	6.98	12.33
MMLU-college computer science	2.71	1.30	0.45	1.54	0.97	31.30	16.90	13.21	9.74	9.79
MMLU-human sexuality	0.71	0.01	0.37	0.55	0.36	35.90	10.32	17.98	19.11	7.53
MMLU-formal logic	1.31	0.79	0.91	0.42	0.81	10.65	5.89	7.43	5.44	7.25
GSM8K	0.31	0.07	0.07	0.03	0.14	27.30	13.79	18.98	14.31	10.90
MATH	0.12	0.03	0.05	0.05	0.06	21.43	5.44	9.76	8.94	8.17
MBPP	0.10	0.00	0.00	0.00	0.00	6.20	2.30	3.60	2.80	1.20

As shown in Table 1 Hate Speech biases notably deteriorated Llama3’s performance in college computer science and human sexuality. Biases towards disabled individuals, as depicted by CrowS Pairs, universally degraded performance across all knowledge-based Q&A tasks, indicating a negative bias association within Llama3’s deeper layers. Gender-related biases did not significantly affect performance. National biases prominently impacted outcomes in college computer science and formal logic, suggesting stereotypical assumptions about educational and professional attributes based on nationality. Appearance-related biases predominantly influenced human sexuality performance.

Knowledge-based Q&A tasks were generally more vulnerable to implicit biases, whereas reasoning tasks such as GSM8K, MATH, and MBPP appeared largely immune, likely due to the nature of reasoning problems that resists bias introduction via RLHF. Interestingly, MBPP’s performance was unaffected by biases that significantly impaired results in computer science, an observation that, according to alignment theory (Contributors, 2023), suggests a decoupling of ‘computer knowledge’ and ‘programming skills’ within LLMs. Our BTBR effectively reduced the detrimental impacts of implicit biases across diverse tasks, as summarized in Table 1.

4.5 Ablation Studies

One might wonder, *why not simply extract the entire bias dataset from LLMs? Are Bayesian methods for data filtering truly necessary?* We address this question by showcasing the effects of over-removal of data in this section. Table 2 compares

AMSD performance between partial data removal using BTBR and complete bias dataset extraction. While BTBR incurred minimal performance losses compared to the baseline Llama3, completely removing a bias dataset led to substantial declines, particularly with Hate Speech where most content represents general knowledge rather than bias. Such variability across datasets highlights the precision of our log-likelihood differential approach in gauging bias extent, where a higher differential denotes a stronger capture of bias by LLMs and a lower one indicates predominant commonsense content.

5 Conclusion

In this research, we conducted an extensive examination of implicit biases within LLMs and introduced a novel approach to mitigate this issue. To address the implicit bias issues, we developed a framework, named BTBR, that employs Bayesian inference techniques to accurately detect and eliminate biases using publicly available datasets. Moreover, we introduced multiple evaluation metrics with diverse evaluation datasets to thoroughly evaluate the LLMs’ performance and fairness after mitigating biases. The results demonstrate that the BTBR framework significantly enhances the fairness of LLMs while preserving high levels of task performance. Not only does this finding validate the efficacy of our methodology, but it also offers fresh perspectives and methodologies for addressing bias in future LLM research and applications.

Limitations

While our primary focus is on addressing the implicit bias in Language Models (LLM), we anticipate broader applications of the BTBR framework in various aspects of LLM fairness. However, enhancing the fairness of LLM is a formidable, long-term task. Achieving the optimal solution may necessitate collaborative efforts across multiple academic and practical fields (Shumailov et al., 2024; Eloundou et al., 2024). Specifically, our implementation of BTBR relies on publicly available datasets to infer hidden biases in LLM, but these datasets may have inherent limitations. For instance, a dataset of hate speech from white supremacist forums might not encompass all types of hate speech biases. Thus, the effectiveness of our bias mitigation strategy is directly tied to the quality of these datasets, underscoring the need for high-quality data sources. Furthermore, translating biased expressions into subject-relation-object triples might not fully grasp the complexity of linguistic bias. Moreover, some biases may span across long contexts, presenting challenges to identification methods based on likelihood ratios. Future plans will aim to extend BTBR to address a broader range of bias types.

References

Jaimeen Ahn and Alice Oh. 2021. Mitigating language-dependent ethnic bias in bert. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 533–549.

Richard D Ashmore and Frances K Del Boca. 1979. Sex stereotypes and implicit personality theory: Toward a cognitive—social psychological conceptualization. *Sex roles*, 5(2):219–248.

Jacob Austin, Augustus Odena, Maxwell Nye, Maarten Bosma, Henryk Michalewski, David Dohan, Ellen Jiang, Carrie Cai, Michael Terry, Quoc Le, et al. 2021. *Program synthesis with large language models*. *ArXiv preprint*, abs/2108.07732.

Rajas Bansal. 2022. *A survey on bias and fairness in natural language processing*. *ArXiv preprint*, abs/2204.09591.

Soumya Barikeri, Anne Lauscher, Ivan Vulić, and Goran Glavaš. 2021. Redditbias: A real-world resource for bias evaluation and debiasing of conversational language models. In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 1941–1955.

Su Lin Blodgett, Solon Barocas, Hal Daumé III, and Hanna Wallach. 2020. *Language (technology) is power: A critical survey of “bias” in NLP*. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 5454–5476, Online. Association for Computational Linguistics.

Su Lin Blodgett, Gilsinia Lopez, Alexandra Olteanu, Robert Sim, and Hanna Wallach. 2021. *Stereotyping Norwegian salmon: An inventory of pitfalls in fairness benchmark datasets*. In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 1004–1015, Online. Association for Computational Linguistics.

Aylin Caliskan, Joanna J Bryson, and Arvind Narayanan. 2017. Semantics derived automatically from language corpora contain human-like biases. *Science*, 356(6334):183–186.

Daniel Cer, Mona Diab, Eneko Agirre, Iñigo Lopez-Gazpio, and Lucia Specia. 2017. *SemEval-2017 task 1: Semantic textual similarity multilingual and crosslingual focused evaluation*. In *Proceedings of the 11th International Workshop on Semantic Evaluation (SemEval-2017)*, pages 1–14, Vancouver, Canada. Association for Computational Linguistics.

Hyeong Kyu Choi and Yixuan Li. 2024. *PICLe: Eliciting diverse behaviors from large language models with persona in-context learning*. In *Forty-first International Conference on Machine Learning*.

Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John Schulman. 2021. *Training verifiers to solve math word problems*. *ArXiv preprint*, abs/2110.14168.

OpenCompass Contributors. 2023. *Opencompass: A universal evaluation platform for foundation models*. <https://github.com/open-compass/opencompass>.

Ona de Gibert, Naiara Perez, Aitor García-Pablos, and Montse Cuadros. 2018. *Hate speech dataset from a white supremacy forum*. In *Proceedings of the 2nd Workshop on Abusive Language Online (ALW2)*, pages 11–20, Brussels, Belgium. Association for Computational Linguistics.

Pieter Delobelle and Bettina Berendt. 2022. Fairdistillation: mitigating stereotyping in language models. In *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*, pages 638–654. Springer.

Jwala Dhamala, Tony Sun, Varun Kumar, Satyapriya Krishna, Yada Pruksachatkun, Kai-Wei Chang, and Rahul Gupta. 2021. Bold: Dataset and metrics for measuring biases in open-ended language generation. In *Proceedings of the 2021 ACM conference*

739	on fairness, accountability, and transparency, pages	9th International Conference on Learning Representations, ICLR 2021, Virtual Event, Austria, May 3-7, 2021. OpenReview.net.	793
740	862–872.		794
741	Peng Ding, Jun Kuang, Dan Ma, Xuezhi Cao, Yun-		795
742	sen Xian, Jiajun Chen, and Shujian Huang. 2023.	Dan Hendrycks, Collin Burns, Steven Basart, Andy	796
743	A wolf in sheep’s clothing: Generalized nested jail-	Zou, Mantas Mazeika, Dawn Song, and Jacob Stein-	797
744	break prompts can fool large language models easily.	hardt. 2021b. Measuring massive multitask language	798
745	<i>Preprint</i> , arXiv:2311.08268.	understanding . In <i>9th International Conference on</i>	799
746	Naomi Ellemers. 2018. Gender stereotypes. <i>Annual</i>	<i>Learning Representations, ICLR 2021, Virtual Event,</i>	800
747	<i>review of psychology</i> , 69(1):275–298.	<i>Austria, May 3-7, 2021</i> . OpenReview.net.	801
748	Tyna Eloundou, Alex Beutel, David G Robinson, Keren	Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul	802
749	Gu-Lemberg, Anna-Luisa Brakman, Pamela Mishkin,	Arora, Steven Basart, Eric Tang, Dawn Song, and	803
750	Meghan Shah, Johannes Heidecke, Lilian Weng, and	Jacob Steinhardt. 2021c. Measuring mathematical	804
751	Adam Tauman Kalai. 2024. First-person fairness in	problem solving with the math dataset. <i>NeurIPS</i> .	805
752	chatbots. <i>arXiv preprint arXiv:2410.19803</i> .		
753	Nikhil Garg, Londa Schiebinger, Dan Jurafsky, and	James L Hilton and William Von Hippel. 1996. Stereo-	806
754	James Zou. 2018. Word embeddings quantify 100	types. <i>Annual review of psychology</i> , 47(1):237–271.	807
755	years of gender and ethnic stereotypes. <i>Proceedings</i>		
756	<i>of the National Academy of Sciences</i> , 115(16):E3635–	Yihuai Hong, Lei Yu, Shauli Ravfogel, Haiqin Yang,	808
757	E3644.	and Mor Geva. 2024. Intrinsic evaluation of un-	809
		learning using parametric knowledge traces . <i>ArXiv</i>	810
		<i>preprint</i> , abs/2406.11614.	811
758	Micah Goldblum, Steven Reich, Liam Fowl, Renkun Ni,	Po-Sen Huang, Huan Zhang, Ray Jiang, Robert Stan-	812
759	Valeria Cherepanova, and Tom Goldstein. 2020. Un-	forth, Johannes Welbl, Jack Rae, Vishal Maini, Dani	813
760	raveling meta-learning: Understanding feature repre-	Yogatama, and Pushmeet Kohli. 2020. Reducing sen-	814
761	sentations for few-shot tasks . In <i>Proceedings of the</i>	timent bias in language models via counterfactual	815
762	<i>37th International Conference on Machine Learning,</i>	evaluation. In <i>Findings of the Association for Com-</i>	816
763	<i>ICML 2020, 13-18 July 2020, Virtual Event</i> , volume	<i>putational Linguistics: EMNLP 2020</i> , pages 65–83.	817
764	119 of <i>Proceedings of Machine Learning Research</i> ,		
765	pages 3607–3616. PMLR.	Ray Jiang, Aldo Pacchiano, Tom Stepleton, Heinrich	818
766	Akshat Gupta and Gopala Anumanchipalli. 2024. Re-	Jiang, and Silvia Chiappa. 2019. Wasserstein fair	819
767	building rome: Resolving model collapse dur-	classification . In <i>Proceedings of the Thirty-Fifth Con-</i>	820
768	ing sequential model editing. <i>arXiv preprint</i>	<i>ference on Uncertainty in Artificial Intelligence, UAI</i>	821
769	<i>arXiv:2403.07175</i> .	2019, Tel Aviv, Israel, July 22-25, 2019, volume 115	822
770	Akshat Gupta, Anurag Rao, and Gopala Anu-	of <i>Proceedings of Machine Learning Research</i> , pages	823
771	manchipalli. 2024a. Model editing at scale leads	862–872. AUAI Press.	824
772	to gradual and catastrophic forgetting. <i>arXiv preprint</i>		
773	<i>arXiv:2401.07453</i> .	Sachin Kumar, Vidhisha Balachandran, Lucille Njoo,	825
774	Akshat Gupta, Dev Sajani, and Gopala Anu-	Antonios Anastasopoulos, and Yulia Tsvetkov. 2023.	826
775	manchipalli. 2024b. A unified framework for model	Language generation models can cause harm: So	827
776	editing . <i>ArXiv preprint</i> , abs/2403.14236.	what can we do about it? an actionable survey . In	828
777	Judith A Hall and Jin X Goh. 2017. Studying stereo-	<i>Proceedings of the 17th Conference of the European</i>	829
778	type accuracy from an integrative social-personality	<i>Chapter of the Association for Computational Lin-</i>	830
779	perspective. <i>Social and Personality Psychology Com-</i>	<i>guistics</i> , pages 3299–3321, Dubrovnik, Croatia. As-	831
780	<i>pass</i> , 11(11):e12357.	sociation for Computational Linguistics.	832
781	Peter Hase, Mohit Bansal, Been Kim, and Asma Ghan-	Meghdad Kurmanji, Peter Triantafillou, Jamie Hayes,	833
782	deharioun. 2024. Does localization inform editing?	and Eleni Triantafillou. 2024. Towards unbounded	834
783	surprising differences in causality-based localization	machine unlearning. <i>Advances in neural information</i>	835
784	vs. knowledge editing in language models. <i>Advances</i>	<i>processing systems</i> , 36.	836
785	<i>in Neural Information Processing Systems</i> , 36.		
786	Pengcheng He, Jianfeng Gao, and Weizhu Chen. 2021.	Guokun Lai, Qizhe Xie, Hanxiao Liu, Yiming Yang,	837
787	Debertav3: Improving deberta using electra-style pre-	and Eduard Hovy. 2017. Race: Large-scale read-	838
788	training with gradient-disentangled embedding shar-	ing comprehension dataset from examinations. In	839
789	ing. <i>arXiv preprint arXiv:2111.09543</i> .	<i>Proceedings of the 2017 Conference on Empirical</i>	840
790	Dan Hendrycks, Collin Burns, Steven Basart, Andrew	<i>Methods in Natural Language Processing</i> , pages 785–	841
791	Critch, Jerry Li, Dawn Song, and Jacob Steinhardt.	794.	842
792	2021a. Aligning AI with shared human values . In	Anne Lauscher, Tobias Lueken, and Goran Glavaš. 2021.	843
		Sustainable modular debiasing of language models.	844
		In <i>Findings of the Association for Computational</i>	845
		<i>Linguistics: EMNLP 2021</i> , pages 4782–4797.	846

847	Hector Levesque, Ernest Davis, and Leora Morgenstern.	Kevin Meng, David Bau, Alex Andonian, and Yonatan	900
848	2012. The winograd schema challenge. In <i>Thir-</i>	Belinkov. 2022a. Locating and editing factual asso-	901
849	<i>teenth international conference on the principles of</i>	ciations in GPT. <i>Advances in Neural Information</i>	902
850	<i>knowledge representation and reasoning</i> .	<i>Processing Systems</i> , 35.	903
851	Shahar Levy, Koren Lazar, and Gabriel Stanovsky. 2021.	Kevin Meng, Arnab Sen Sharma, Alex Andonian,	904
852	Collecting a large-scale gender bias dataset for coref-	Yonatan Belinkov, and David Bau. 2022b. Mass	905
853	erence resolution and machine translation. In <i>Find-</i>	editing memory in a transformer . <i>ArXiv preprint</i> ,	906
854	<i>ings of the Association for Computational Linguistics:</i>	abs/2210.07229 .	907
855	<i>EMNLP 2021</i> , pages 2470–2480.		
856	Tao Li, Daniel Khashabi, Tushar Khot, Ashish Sabhar-	Kevin Meng, Arnab Sen Sharma, Alex Andonian,	908
857	wal, and Vivek Srikumar. 2020. Uncovering stereo-	Yonatan Belinkov, and David Bau. 2022c. Mass-	909
858	typing biases via underspecified questions. In <i>Find-</i>	editing memory in a transformer . <i>ArXiv preprint</i> ,	910
859	<i>ings of the Association for Computational Linguistics:</i>	abs/2210.07229 .	911
860	<i>EMNLP 2020</i> , pages 3475–3489.		
861	Yingji Li, Mengnan Du, Rui Song, Xin Wang, and Ying	AI Meta. 2024. Introducing meta llama 3: The most	912
862	Wang. 2023. A survey on fairness in large language	capable openly available llm to date. <i>Meta AI</i> .	913
863	models . <i>ArXiv preprint</i> , abs/2308.10149 .		
864	Yong Lin, Lu Tan, Hangyu Lin, Zeming Zheng, Renjie	Moin Nadeem, Anna Bethke, and Siva Reddy. 2021.	914
865	Pi, Jipeng Zhang, Shizhe Diao, Haoxiang Wang, Han	Stereoset: Measuring stereotypical bias in pretrained	915
866	Zhao, Yuan Yao, et al. 2023. Speciality vs gener-	language models. In <i>Proceedings of the 59th Annual</i>	916
867	ality: An empirical study on catastrophic forgetting	<i>Meeting of the Association for Computational Lin-</i>	917
868	in fine-tuning foundation models . <i>ArXiv preprint</i> ,	<i>guistics and the 11th International Joint Conference</i>	918
869	abs/2309.06256 .	<i>on Natural Language Processing (Volume 1: Long</i>	919
		<i>Papers)</i> , pages 5356–5371.	920
870	Zhisheng Lin, Han Fu, Chenghao Liu, Zhuo Li, and Jian-	Nikita Nangia, Clara Vania, Rasika Bhalerao, and	921
871	ling Sun. 2024. Pemt: Multi-task correlation guided	Samuel R. Bowman. 2020. CrowS-pairs: A chal-	922
872	mixture-of-experts enables parameter-efficient trans-	lenge dataset for measuring social biases in masked	923
873	fer learning. <i>arXiv preprint arXiv:2402.15082</i> .	language models . In <i>Proceedings of the 2020 Con-</i>	924
		<i>ference on Empirical Methods in Natural Language</i>	925
874	Qijun Luo, Hengxu Yu, and Xiao Li. 2024. Badam:	<i>Processing (EMNLP)</i> , pages 1953–1967, Online. As-	926
875	A memory efficient full parameter training method	sociation for Computational Linguistics.	927
876	for large language models . <i>ArXiv preprint</i> ,		
877	abs/2404.02827 .	Debora Nozza, Federico Bianchi, Dirk Hovy, et al. 2021.	928
878	Aengus Lynch, Phillip Guo, Aidan Ewart, Stephen	Honest: Measuring hurtful sentence completion in	929
879	Casper, and Dylan Hadfield-Menell. 2024. Eight	language models. In <i>Proceedings of the 2021 Con-</i>	930
880	methods to evaluate robust unlearning in llms . <i>ArXiv</i>	<i>ference of the North American Chapter of the Asso-</i>	931
881	<i>preprint</i> , abs/2402.16835 .	<i>ciation for Computational Linguistics: Human Lan-</i>	932
882	Xinyao Ma, Maarten Sap, Hannah Rashkin, and Yejin	<i>guage Technologies</i> . Association for Computational	933
883	Choi. 2020. PowerTransformer: Unsupervised con-	Linguistics.	934
884	trollable revision for biased language correction . In		
885	<i>Proceedings of the 2020 Conference on Empirical</i>	Kartikey Pant and Tanvi Dadu. 2022. Incorporating	935
886	<i>Methods in Natural Language Processing (EMNLP)</i> ,	subjectivity into gendered ambiguous pronoun (gap)	936
887	pages 7426–7441, Online. Association for Computa-	resolution using style transfer. In <i>Proceedings of the</i>	937
888	tional Linguistics.	<i>4th Workshop on Gender Bias in Natural Language</i>	938
889	Chandler May, Alex Wang, Shikha Bordia, Samuel Bow-	<i>Processing (GeBNLP)</i> , pages 273–281.	939
890	man, and Rachel Rudinger. 2019. On measuring so-		
891	cial biases in sentence encoders. In <i>Proceedings of</i>	Alicia Parrish, Angelica Chen, Nikita Nangia,	940
892	<i>the 2019 Conference of the North American Chap-</i>	Vishakh Padmakumar, Jason Phang, Jana Thompson,	941
893	<i>ter of the Association for Computational Linguistics:</i>	Phu Mon Htut, and Samuel R Bowman. 2021. Bbq:	942
894	<i>Human Language Technologies, Volume 1 (Long and</i>	A hand-built bias benchmark for question answering.	943
895	<i>Short Papers)</i> , pages 622–628.	<i>arXiv preprint arXiv:2110.08193</i> .	944
896	Ninareh Mehrabi, Fred Morstatter, Nripsuta Saxena,	Vaidehi Patil, Peter Hase, and Mohit Bansal. 2023. Can	945
897	Kristina Lerman, and Aram Galstyan. 2021. A sur-	sensitive information be deleted from llms? objec-	946
898	vey on bias and fairness in machine learning. <i>ACM</i>	tives for defending against extraction attacks . <i>ArXiv</i>	947
899	<i>computing surveys (CSUR)</i> , 54(6):1–35.	<i>preprint</i> , abs/2309.17410 .	948
		Cheryl Pritlove, Clara Juando-Prats, Kari Ala-	949
		Leppilampi, and Janet A Parsons. 2019. The good,	950
		the bad, and the ugly of implicit bias. <i>The Lancet</i> ,	951
		393(10171):502–504.	952

953	Rebecca Qian, Candace Ross, Jude Fernandes,	57th Annual Meeting of the Association for Computa-	1010
954	Eric Michael Smith, Douwe Kiela, and Adina	tional Linguistics, pages 1630–1640, Florence, Italy.	1011
955	Williams. 2022. Perturbation augmentation for fairer	Association for Computational Linguistics.	1012
956	NLP . In <i>Proceedings of the 2022 Conference on</i>		
957	<i>Empirical Methods in Natural Language Processing</i> ,	Xiaoyu Tan, Yongxin Deng, Xihe Qiu, Weidi Xu, Chao	1013
958	pages 9496–9521, Abu Dhabi, United Arab Emirates.	Qu, Wei Chu, Yinghui Xu, and Yuan Qi. 2024.	1014
959	Association for Computational Linguistics.	Thought-like-pro: Enhancing reasoning of large lan-	1015
		guage models through self-driven prolog-based chain-	1016
960	David Rein, Betty Li Hou, Asa Cooper Stickland,	of-thought . <i>ArXiv preprint</i> , abs/2407.14562.	1017
961	Jackson Petty, Richard Yuanzhe Pang, Julien Di-		
962	rani, Julian Michael, and Samuel R. Bowman. 2023.	Yi Chern Tan and L Elisa Celis. 2019. Assessing so-	1018
963	Gpqa: A graduate-level google-proof q&a bench-	cial and intersectional biases in contextualized word	1019
964	mark . <i>Preprint</i> , arXiv:2311.12022.	representations. <i>Advances in neural information pro-</i>	1020
		<i>cessing systems</i> , 32.	1021
965	Rachel Rudinger, Jason Naradowsky, Brian Leonard,		
966	and Benjamin Van Durme. 2018. Gender bias in	Eva Vanmassenhove, Chris Emmery, and Dimitar Shteri-	1022
967	coreference resolution. In <i>Proceedings of the 2018</i>	onov. 2021. NeuTral Rewriter: A rule-based and neu-	1023
968	<i>Conference of the North American Chapter of the</i>	ral approach to automatic rewriting into gender neu-	1024
969	<i>Association for Computational Linguistics: Human</i>	tral alternatives . In <i>Proceedings of the 2021 Confer-</i>	1025
970	<i>Language Technologies, Volume 2 (Short Papers)</i> ,	<i>ence on Empirical Methods in Natural Language Pro-</i>	1026
971	pages 8–14.	<i>cessing</i> , pages 8940–8948, Online and Punta Cana,	1027
972	Leonard Salewski, Stephan Alaniz, Isabel Rio-Torto,	Dominican Republic. Association for Computational	1028
973	Eric Schulz, and Zeynep Akata. 2024. In-context im-	Linguistics.	1029
974	personation reveals large language models’ strengths		
975	and biases. <i>Advances in Neural Information Process-</i>	Mengru Wang, Ningyu Zhang, Ziwen Xu, Zekun Xi,	1030
976	<i>ing Systems</i> , 36.	Shumin Deng, Yunzhi Yao, Qishen Zhang, Linyi	1031
977	Emily Sheng, Kai-Wei Chang, Premkumar Natarajan,	Yang, Jindong Wang, and Huajun Chen. 2024. Detox-	1032
978	and Nanyun Peng. 2019. The woman worked as	ifying large language models via knowledge editing.	1033
979	a babysitter: On biases in language generation . In	<i>arXiv preprint arXiv:2403.14472</i> .	1034
980	<i>Proceedings of the 2019 Conference on Empirical</i>	Alexander Warnecke, Lukas Pirch, Christian Wressneg-	1035
981	<i>Methods in Natural Language Processing and the</i>	ger, and Konrad Rieck. 2021. Machine unlearning of	1036
982	<i>9th International Joint Conference on Natural Lan-</i>	features and labels . <i>ArXiv preprint</i> , abs/2108.11577.	1037
983	<i>guage Processing (EMNLP-IJCNLP)</i> , pages 3407–		
984	3412, Hong Kong, China. Association for Computa-	Kellie Webster, Marta Recasens, Vera Axelrod, and Ja-	1038
985	tional Linguistics.	son Baldrige. 2018. Mind the gap: A balanced	1039
986	Ilia Shumailov, Jamie Hayes, Eleni Triantafillou,	corpus of gendered ambiguous pronouns. <i>Transac-</i>	1040
987	Guillermo Ortiz-Jimenez, Nicolas Papernot, Matthew	<i>tions of the Association for Computational Linguis-</i>	1041
988	Jagielski, Itay Yona, Heidi Howard, and Eugene Bag-	<i>tics</i> , 6:605–617.	1042
989	dasaryan. 2024. Ununlearning: Unlearning is not	Kellie Webster, Xuezhi Wang, Ian Tenney, Alex Beutel,	1043
990	sufficient for content regulation in advanced genera-	Emily Pitler, Ellie Pavlick, Jilin Chen, Ed Chi, and	1044
991	tive ai . <i>ArXiv preprint</i> , abs/2407.00106.	Slav Petrov. 2020. Measuring and reducing gendered	1045
992	Eric Michael Smith, Melissa Hall, Melanie Kambadur,	correlations in pre-trained models. <i>arXiv preprint</i>	1046
993	Eleonora Presani, and Adina Williams. 2022. “i’m	<i>arXiv:2010.06032</i> .	1047
994	sorry to hear that”: Finding new biases in language	Yotam Wolf, Noam Wies, Oshri Avnery, Yoav Levine,	1048
995	models with a holistic descriptor dataset. In <i>Proceed-</i>	and Amnon Shashua. 2023. Fundamental limita-	1049
996	<i>ings of the 2022 Conference on Empirical Methods</i>	tions of alignment in large language models . <i>ArXiv</i>	1050
997	<i>in Natural Language Processing</i> , pages 9180–9211.	<i>preprint</i> , abs/2304.11082.	1051
998	Niklas Stoehr, Mitchell Gordon, Chiyuan Zhang, and	Sang Michael Xie, Aditi Raghunathan, Percy Liang,	1052
999	Owen Lewis. 2024. Localizing paragraph mem-	and Tengyu Ma. 2022. An explanation of in-context	1053
1000	orization in language models . <i>ArXiv preprint</i> ,	learning as implicit bayesian inference . In <i>The Tenth</i>	1054
1001	abs/2403.19851.	<i>International Conference on Learning Representa-</i>	1055
1002	Natalie Jomini Stroud. 2008. Media use and political	<i>tions, ICLR 2022, Virtual Event, April 25-29, 2022</i> .	1056
1003	predispositions: Revisiting the concept of selective	OpenReview.net.	1057
1004	exposure. <i>Political behavior</i> , 30:341–366.	Zhongbin Xie and Thomas Lukasiewicz. 2023. An em-	1058
1005	Tony Sun, Andrew Gaut, Shirlyn Tang, Yuxin Huang,	pirical analysis of parameter-efficient methods for de-	1059
1006	Mai ElSherief, Jieyu Zhao, Diba Mirza, Elizabeth	biasing pre-trained language models . <i>ArXiv preprint</i> ,	1060
1007	Belding, Kai-Wei Chang, and William Yang Wang.	abs/2306.04067.	1061
1008	2019. Mitigating gender bias in natural language	Junsang Yoon, Akshat Gupta, and Gopala Anu-	1062
1009	processing: Literature review . In <i>Proceedings of the</i>	manchipalli . 2024. Is bigger edit batch size always	1063
		better?—an empirical study on model editing with	1064
		llama-3 . <i>ArXiv preprint</i> , abs/2405.00664.	1065

Ge Zhang, Scott Qu, Jiaheng Liu, Chenchen Zhang, Chenghua Lin, Chou Leuang Yu, Danny Pan, Esther Cheng, Jie Liu, Qunshu Lin, Raven Yuan, Tune Zheng, Wei Pang, Xinrun Du, Yiming Liang, Yinghao Ma, Yizhi Li, Ziyang Ma, Bill Lin, Emmanouil Benetos, Huan Yang, Junting Zhou, Kaijing Ma, Minghao Liu, Morry Niu, Noah Wang, Quehry Que, Ruibo Liu, Sine Liu, Shawn Guo, Soren Gao, Wangchunshu Zhou, Xinyue Zhang, Yizhi Zhou, Yubo Wang, Yuelin Bai, Yuhang Zhang, Yuxiang Zhang, Zenith Wang, Zhenzhu Yang, Zijian Zhao, Jiajun Zhang, Wanli Ouyang, Wenhao Huang, and Wenhui Chen. 2024. Map-neo: Highly capable and transparent bilingual large language model series. *arXiv preprint arXiv: 2405.19327*.

Jieyu Zhao, Tianlu Wang, Mark Yatskar, Vicente Ordonez, and Kai-Wei Chang. 2018. *Gender bias in coreference resolution: Evaluation and debiasing methods*. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 2 (Short Papers)*, pages 15–20, New Orleans, Louisiana. Association for Computational Linguistics.

Kairan Zhao, Meghdad Kurmanji, George-Octavian Bărbulescu, Eleni Triantafyllou, and Peter Triantafyllou. 2024. *What makes unlearning hard and what to do about it*. *ArXiv preprint*, abs/2406.01257.

Andy Zou, Zifan Wang, Nicholas Carlini, Milad Nasr, J Zico Kolter, and Matt Fredrikson. 2023. *Universal and transferable adversarial attacks on aligned language models*. *ArXiv preprint*, abs/2307.15043.

A Related Work

The issue of fairness in LLMs is commonly evaluated using two main types of indicators (Li et al., 2023): intrinsic bias evaluation indicators and extrinsic bias evaluation indicators. Intrinsic bias evaluation indicators are applied in embeddings and formalize intrinsic bias by statistically quantifying the association between targets and concepts. These indicators include similarity-based measurement methods (Caliskan et al., 2017; May et al., 2019; Lauscher et al., 2021; Tan and Celis, 2019) and probability-based indicators (Webster et al., 2020; Ahn and Oh, 2021; Nadeem et al., 2021). For extrinsic bias evaluation, benchmarks like WinoBias (Zhao et al., 2018) and Winogender (Rudinger et al., 2018), developed based on Winograd (Levesque et al., 2012) format datasets, are used for coreference resolution. An LLM is considered to exhibit gender bias if it more accurately associates a pronoun with a profession dominated by that pronoun’s gender as opposed to a non-dominated profession. The bias score is defined as the difference between the model’s evaluation of

“stereotypes” and “counter-stereotypes”. Various extensions (Levy et al., 2021; Webster et al., 2018; Pant and Dadu, 2022) have been proposed to create different coreference resolution evaluation datasets based on WinoBias and Winogender.

Furthermore, some studies consider the semantic similarity between sentence pairs to evaluate the association between gender and profession. For instance, (Webster et al., 2020) collected 276 sentences from STS-B (Cer et al., 2017) and created a series of neutral sentence pair templates, one sentence containing gender terms and the other containing a profession with gender implications, such as “A [woman] is walking.” and “A [nurse] is walking.”. A gender-neutral model should provide the same similarity estimates for these two sets of gender sentence pairs, with any difference indicating gender bias.

Certain studies (Dhamala et al., 2021; Sheng et al., 2019; Huang et al., 2020; Jiang et al., 2019; Nozza et al., 2021; Smith et al., 2022; Barikeri et al., 2021; Parrish et al., 2021; He et al., 2021; Lai et al., 2017; Li et al., 2020) fine-tune evaluation datasets containing different condition prompts and then evaluate the generated content. For example, to measure religious stereotypes, a vague scenario like “A Christian man and a Muslim man are arguing fiercely.” is given, along with the question “Who looks like a criminal?” This approach aims to exclude confusing factors that could bias the model’s measurements by providing a context without a clear answer.

Our work appears to align more closely with external bias evaluation, characterizing external bias by measuring performance gaps. However, unlike studies that directly measure the explicit bias tendencies claimed by the model, our work focuses more on the issue of ‘implicit bias’, which has been overlooked in existing research.

Our work also involves removing bias from models. (Hong et al., 2024) introduced the concept of “knowledge representations”, assessing whether the unlearning algorithm truly erases information from the model weights (instead of concealing them until induced by malicious users) by calculating changes in the concept vectors in LLMs. (Hong et al., 2024)’s experiments revealed that all methods based on fine-tuning essentially have minimal changes to concept vectors. In contrast, knowledge editing methods such as MEMIT (Meng et al., 2022c), PEMT (Lin et al., 2024), DINM (Wang et al., 2024), R-ROME (Gupta and Anumanchipalli,

2024; Gupta et al., 2024a), EMMET(Gupta et al., 2024b,a), and FT-L(Meng et al., 2022a) are capable of disrupting knowledge encoded in LLMs. Given that MEMIT has been extensively researched and validated, we can determine its optimal hyperparameter settings based on existing information and ensure the overall stability of the entire BTBR workflow. This allows us to concentrate on our core objective: identifying the most representative bias data in a dataset and assessing the potential impact on performance if an excessive amount of non-biased data is erroneously removed. We acknowledge that within our BTBR framework, MEMIT could be replaced with any bias removal method.

B Additional experimental details

B.1 Hardware Setup and Hyperparameter Selection

Our experiments were conducted using a single NVIDIA A800-80GB GPU. Regarding hyperparameters, we set the temperature to 0.6 and top_p to 0.9 for any LLM inference involved, following official recommendations for Llama (Meta, 2024). As mentioned in Section 3.1, we used fine-tuned models to estimate $p_{\theta}(\mathbf{x}|\phi')$. To mitigate the computational costs of fine-tuning, we employed BAdam (Luo et al., 2024), an optimization method utilizing the block coordinate descent framework with Adam as the inner solver, treating each transformer layer module as a separate block and training one block at a time. Adhering to BAdam’s official guidelines for Llama3 training, we set the learning rate at $1e-6$, with block switching frequency at every 100 epochs for a total of three epochs. Moreover, from an intuitive perspective, the choice of the hyperparameter K is influenced by the characteristics of the biased dataset; the larger the number of purely biased data points present in the dataset, the greater the value of K should be, and conversely. We have illustrated the DB values for a subset of the Hate Speech dataset in Figure 3. In this instance, we opted for $K = 30$.

To eliminate bias from LLMs, we employed the MEMIT method for model editing. Originally, MEMIT edited multiple LLM layers simultaneously, but findings by (Gupta et al., 2024b) suggested that multi-layer editing could obscure actual editing performance. Therefore, our experiments focused on editing a single layer. (Meng et al., 2022c) evaluated hidden states in LLMs for fact recall through causal tracing; however, later research

(Hase et al., 2024) indicated that layers identified as significant didn’t necessarily correlate with editing performance. Empirically, (Yoon et al., 2024) identified the most effective layer for editing in Llama models (including Llama2 and Llama3); consistently, editing the 1 layer yielded better outcomes, thus, our experiments also targeted this layer. It should be noted that in Llama3-8B, layers are indexed from 0 to 31. Moreover, considering that editing efficacy diminishes with larger batch sizes (Yoon et al., 2024), we opted for sequential editing with a batch size of one.

B.2 Automated Model Editing

In Section 3.2, we discussed our approach to automatically address the identified biases, which involves transforming xx from natural language into structured knowledge. To facilitate this, we employed a few-shot learning technique to guide the LLM in performing the rewriting task. Below, we provide a demonstration of the prompt utilized during this transformation process. In this demonstration, the content preceding the **red hash line** constitutes the input, while the following content represents the LLM’s response. For additional demonstrations, please refer to the supplementary material.

Prompt demonstration

System:

You are a helpful assistant that helps people solve problems.

User:

Your task is to help me extract subject-relation-object triples from sentences. Below are the detailed instructions. Thanks for your help!

****Instructions Start****

- (1) You should carefully read each sentence provided.
- (2) Identify the main subject, the relation (verb or action), and the object in the sentence. The subject is the entity performing the action, the relation is the action itself, and the object is the entity that is receiving the action.
- (3) Your extraction should be clear and explicit, structured as "subject + relation + object".
- (4) Please present each SRO triple in the format: "Subject: XXXX; Relation: XXXX; Object: XXXX".
- (5) After extracting the triple, write it down. It should be in a new line, starting with "####".

Many thanks for your help! I am looking forward to your response!

****Instructions End****

Here is the problem:

The white man has no future in Canada, that is for sure, because the country is being flooded with non-white immigrants.

Assistant:

Sure! I am happy to help you solve this problem. Here is the answer:

Subject: The white man; Relation: has; Object: no future in Canada

Subject: The country; Relation: is being flooded; Object: with non-white immigrants

User:

Excellent work! Here is another problem for you to solve. Please apply the same approach you used for the previous one(s) to tackle this new one.

Problem:

And the sad thing is the white students at those schools will act like that too.

Assistant:

Sure! I am happy to help you solve this problem. Here is the answer:

Subject: the white students at those schools; Relation: will act; Object: like that too

User:

Excellent work! Here is another problem for you to solve. Please apply the same approach you used for the previous one(s) to tackle this new one.

Problem:

They need the white male distracted as much as possible and of course the glorification of the groid; two birds one stone.

Assistant:

Sure! I am happy to help you solve this problem. Here is the answer:

#####

Subject: They; Relation: need; Object: the white male distracted as much as possible

Subject: the glorification; Relation: is; Object: of the groid