

Bias Amplification: Large Language Models as Increasingly Biased Media

Anonymous ACL submission

Abstract

Model collapse—a phenomenon characterized by performance degradation due to iterative training on synthetic data—has been widely studied. However, its implications for bias amplification, the progressive intensification of pre-existing societal biases in Large Language Models (LLMs), remain significantly underexplored, despite the growing influence of LLMs in shaping online discourse. In this paper, we introduce a open, generational, and long-context benchmark specifically designed to measure political bias amplification in LLMs, leveraging sentence continuation tasks derived from a comprehensive dataset of U.S. political news. Our empirical study using GPT-2 reveals consistent and substantial political bias intensification (e.g., right-leaning amplification) over iterative synthetic training cycles. We evaluate three mitigation strategies—Overfitting, Preservation, and Accumulation—and demonstrate that bias amplification persists independently of model collapse, even when the latter is effectively controlled. Furthermore, we propose a mechanistic analysis approach that identifies neurons correlated with specific phenomena during inference through regression and statistical tests. This analysis uncovers largely distinct neuron populations driving bias amplification and model collapse, underscoring fundamentally different underlying mechanisms. Finally, we supplement our empirical findings with theoretical intuition that explains the separate origins of these phenomena, guiding targeted strategies for bias mitigation.

1 Introduction

Large Language Models (LLMs) have become essential tools for content creation and summarization in various sectors, including media, academia, and business (Maslej et al., 2024). However, a significant but underexplored risk arises as LLMs increasingly rely on their own or other synthetic outputs for training, potentially amplifying pre-

existing societal biases (Peña-Fernández et al., 2023; Porlezza and Ferri, 2022; Nishal and Diakopoulos, 2024). This phenomenon, known as bias amplification, refers to the progressive reinforcement and intensification of existing biases through iterative synthetic training (Mehrabi et al., 2022; Taori and Hashimoto, 2022). This issue stems from the inherent tendency of LLMs to learn from biased datasets; prior studies indicate that LLMs readily absorb biases from human-generated text (Parrish et al., 2022; Wang et al., 2024; Bender et al., 2021). Further, fine-tuning LLMs on biased datasets can align them with specific political ideologies (Haller et al., 2023; Rettenberger et al., 2024a). Classifiers trained on synthetic data also increasingly favor specific class labels over time (Wyllie et al., 2024), and diversity tends to decrease, potentially marginalizing certain demographic groups (Alemohammad et al., 2023; Hamilton, 2024). The implications of bias amplification are substantial, including perpetuation of stereotypes, reinforcement of social inequalities, and potential impacts on democratic processes through the skewing of public opinion and increased polarization. Despite its significance, comprehensive frameworks and empirical research specifically addressing bias amplification in language models remain sparse, although related work exists on discriminative models.

In this study, we empirically investigate political bias amplification in GPT-2. We define political bias as the disproportionate generation of content aligned with specific political ideologies. Our experiments reveal that GPT-2 progressively exhibits stronger right-leaning and center-leaning biases in two distinct scenarios: (1) starting from fine-tuning on an unbiased dataset, and (2) initially fine-tuned exclusively on center-leaning articles. We also evaluate three mitigation strategies—Overfitting, Preservation, and Accumulation—to address bias amplification and model collapse. Preservation

043
044
045
046
047
048
049
050
051
052
053
054
055
056
057
058
059
060
061
062
063
064
065
066
067
068
069
070
071
072
073
074
075
076
077
078
079
080
081
082
083

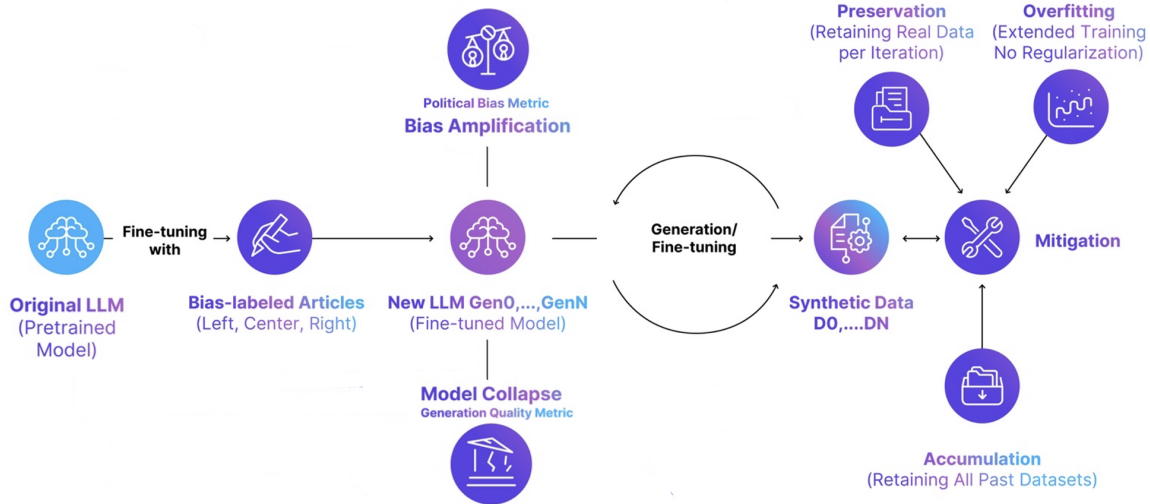


Figure 1: Overview of the iterative experimental procedure for synthetic fine-tuning and analysis.

mitigates both phenomena effectively in the first scenario but fails to prevent bias amplification in the second. Additionally, we propose a novel mechanistic analysis using regression and statistical testing to examine neuron-level changes correlated with bias amplification and model collapse, identifying largely distinct neuron groups for each phenomenon. This suggests different underlying mechanisms for bias amplification and model collapse.

In summary, our contributions are: (i) a highly accurate classifier for detecting political bias in long-text content, providing a benchmark for evaluating political bias in LLMs via sentence continuation tasks; (ii) an empirical assessment of political bias amplification in GPT-2 across two fine-tuning setups; (iii) evaluation of three mitigation strategies; (iv) a novel mechanistic analysis method identifying neurons correlated with specific phenomena during inference; and (v) a theoretical intuition that explains the difference between bias amplification and model collapse based on their underlying causes. The experimental framework presented can be extended to other models and different types of biases.

2 Related Work

Bias Amplification has been studied in various domains. For instance, [Zhao et al. \(2017\)](#) found Conditional Random Fields can worsen social biases from training data, proposing an in-process Lagrangian Relaxation method to align model and data biases. [Mehrabi et al. \(2022\)](#) later described bias amplification in feedback loops, where models amplify existing bias and generate more biased data

through real-world interaction. [Xu et al. \(2023\)](#); [Zhou et al. \(2024\)](#) showed recommendation models amplify mainstream preferences, overrepresenting them and neglecting rarer items, akin to sampling error ([Shumailov et al., 2024](#)).

Classifiers trained on synthetic data increasingly favor certain labels over generations ([Wyllie et al., 2024](#); [Taori and Hashimoto, 2022](#)). Similarly, generative models like Stable Diffusion show bias amplification through feature overrepresentation from training data ([Ferbach et al., 2024](#); [Chen et al., 2024](#)). More recently, [Li et al. \(2025\)](#) investigated gender and cultural bias amplification in LLMs (classification and generation tasks, 1-5 synthetic rounds), proposing pre-processing (labeling bias, removing identity words) and in-processing (penalizing deviation from real data) mitigations; these showed varied effectiveness in one-round fine-tuning.

Model Collapse. Model collapse is a deterioration where models recursively trained on their own output distort reality and lose generalizability (e.g., prioritizing common events, neglecting rare ones, or shifting distributions) ([Shumailov et al., 2024](#); [Alemohammad et al., 2023](#); [Guo et al., 2024](#); [Wyllie et al., 2024](#); [Dohmatob et al., 2024a](#)). [Shumailov et al. \(2024\)](#) showed this with OPT-125M, where perplexity distributions skewed towards lower values with longer tails. Increased repetition in synthetically fine-tuned GPT-2 was noted by [Taori and Hashimoto \(2022\)](#). Performance deterioration in models like OPT-350M, Llama2, and GPT-2 (e.g., reduced linguistic diversity, token probability divergence) after several generations was shown by

Guo et al. (2024); Dohmatob et al. (2024b); Seddik et al. (2024). In generative image models, Alemohammad et al. (2023) found quality and diversity deteriorate with synthetic training; however, user cherry-picking of high-quality outputs (a form of sampling error) helped maintain quality. Hamilton (2024) noted GPT-3.5-turbo exhibited less perspective diversity in narrative writing than earlier models (davinci-instruct-beta, text-davinci-003).

Political Biases. In parallel, growing attention has been paid to political biases in LLMs, now a prevalent form of "media" that people rely on for global news (Maslej et al., 2024). Rettenberger et al. (2024b); Shumailov et al. (2024); Feng et al. (2024) explored the bias through voting simulations within the spectrum of German political parties, consistently finding a left-leaning bias in models like GPT-3 and Llama3-70B. Similarly, for the U.S. political landscape, Rotaru et al. (2024); Motoki et al. (2024) identified a noticeable left-leaning bias in ChatGPT and Gemini when tasked with rating news content, evaluating sources, or responding to political questionnaires. (Bang et al., 2024) study political bias in LLMs through the task of generating news headlines on politically sensitive topics and find that the political perspectives expressed by LLMs vary depending on the subject matter.

3 Methodology

This section provides the details of the experiments on LLMs, focusing on the sequential and synthetic fine-tuning of GPT-2. The step-by-step experimental procedure is outlined in Figure 1. Our study focuses on the political bias of LLMs within the US political spectrum, particularly in sentence continuation tasks. This is important as LLMs are increasingly influencing global news consumption (Maslej et al., 2024; Peña-Fernández et al., 2023; Porlezza and Ferri, 2022), and traditional news outlets, such as the Associated Press, are beginning to integrate LLMs for automated content generation from structured data (The Associated Press, 2024).

3.1 Dataset Preparation

We randomly selected 1,518 articles from the Webis-Bias-Flipper-18 dataset (Chen et al., 2018), which contains political articles from a range of U.S. media outlets published between 2012 and 2018, along with bias ratings assigned at the time for each media source. These bias ratings, provided by AllSides, were determined through a multi-stage

process incorporating assessments from both bipartisan experts and the general public (AllSides, 2024a). The random sampling was stratified based on bias ratings to ensure an even distribution of the 1,518 articles into three groups of 506 each, representing left-leaning, right-leaning, and center-leaning media.

3.2 Successive Fine-tuning

Following Shumailov et al. (2024); Dohmatob et al. (2024b), we perform iterative fine-tuning. First, GPT-2 is fine-tuned on the 1,518 real news articles (detailed in Section 3.1) to yield the Generation 0 (G0) model. G0 then generates a synthetic dataset, D_0 , of the same size (1,518 articles). This dataset D_0 is used to fine-tune the Generation 1 (G1) model, which is the first model trained on purely synthetic data. The process continues up to Generation 10 (G10), where each G_i model is fine-tuned on the synthetic data D_{i-1} produced by model G_{i-1} .

The fine-tuning procedure remained consistent across all experiments unless stated otherwise. The input length was capped at 512 tokens, with the EOS token used for padding. The model was trained for 5 epochs, using a batch size of 8, a learning rate of 5×10^{-5} , and a weight decay of 0.01. Fine-tuning was conducted using standard functionalities available in transformer libraries. After each cycle, the model was saved and used to generate synthetic data for the subsequent iteration.¹

3.3 Synthetic Data Generation

Synthetic datasets, $\{D_i\}_{i=0}^{10}$, are generated as follows: For each original news article, its tokenized title serves as an initial prompt, and its tokenized body is segmented into sequential 64-token blocks, which serve as subsequent prompts. For each such prompt, the model predicts the next 64 tokens. These predictions are made using deterministic generation to enhance the reproducibility. All the newly generated 64-token sequences (one from the title prompt and one from each body block prompt) are concatenated and then decoded back into text. This process creates one synthetic article from each original article, resulting in a synthetic dataset of the same number of articles as the original.

3.4 Political Bias Metric

Metric. We develop a classification model to assess the political leaning of each LLM based on

¹Fine-tuned models will be made public upon acceptance.

its generated synthetic news articles. The model is trained on the Webis-Bias-Flipper-18 dataset, excluding the 1,518 articles used for GPT-2 fine-tuning. To mitigate class imbalance, center-leaning articles are resampled to ensure equal representation across categories. The dataset is then divided into training (70%), validation (15%), and test (15%) subsets, stratified by bias label. Additionally, we conduct a human review to remove any identifiable information about media sources and authors.

Specifically, the training dataset comprises 2,781 distinct events that occurred in the U.S. between 2012 and 2018. For each event, it includes news articles collected from a wide range of media outlets. In total, it contains articles from 97 different outlets, such as The Washington Examiner, The Washington Post, HuffPost, Reuters, and others. This makes the dataset a strong representation of the diversity of U.S. media sources and event domains, and therefore strengthens the generalizability of our classifier in the U.S. context.

After performing a grid search across multiple models, we find that roberta-base achieves the best performance, with an evaluation loss of 0.4035 and a macro F1 score of 0.9196 on the test set (see Table 1). Thus, we select roberta-base as the benchmark for political bias detection in subsequent experiments. It is important to note that this classifier is trained on data from 2012-2018, and its direct application to significantly later content might eventually require recalibration due to the evolving nature of the political landscape and discourse. Details on model training are provided in Appendix A.

Bias Construct. We define political bias in the news article generation task as the disproportionate production of articles with specific political leanings, as identified by our classifier. Unlike Bang et al. (2024), which defines political bias in a topic- or issue-specific manner, we take a broader perspective by measuring the overall political leaning of the model across a diverse set of topics. This approach uses articles from U.S. media outlets published between 2012 and 2018. We use this definition because we view LLMs as analogous to media outlets: while an outlet may publish content spanning the political spectrum, both the public and bipartisan organizations like AllSides assign it an overall political rating based on the average ideological slant of its content and its perceived alignment.

Table 1: Macro F1 Scores for Political Bias Classifier Models; roberta-base selected.

Model	Macro F1 Score
distilbert-base-uncased	0.8308
bert-base-uncased	0.8559
albert-base-v2	0.8649
roberta-base	0.9196

3.5 Generation Quality Metric

Metric. We introduce a metric, *text quality index*, to evaluate generation quality, specifically addressing the issue of repetitive content in later model iterations, which can distort traditional perplexity metrics (see Section 4.2). This metric is based on the Gibberish Detector (Jindal, 2021), which identifies incoherent or nonsensical text. The detector categorizes text into four levels: (1) Noise—individual words hold no meaning, (2) Word Salad—incoherent phrases, (3) Mild Gibberish—grammatical or syntactical distortions, and (4) Clean—coherent, meaningful sentences. To quantify generation quality, each sentence receives a Gibberish score: 3 for Clean, 2 for Mild Gibberish, 1 for Word Salad, and 0 for Noise. The *text quality index* is computed as the average score across all sentences in an article. This metric prioritizes coherence and meaning, offering a more meaningful assessment of generation quality than perplexity.

Motivation. We adopt the gibberish detector as our primary tool for evaluating generation quality because it directly captures a fundamental aspect of model deterioration: the loss of coherence and semantic clarity, which becomes particularly pronounced in later generations of synthetic training as observed in our experiments (see Section 4). While other forms of degradation—such as reduced factual consistency or increased hallucination—may also occur, they are often secondary to the loss of basic intelligibility. As such, we leave the investigation of factual consistency deterioration to future work, particularly in contexts more narrowly focused on diagnosing model collapse.

3.6 Mechanistic Analysis

To gain a clearer understanding of the causes of bias amplification and how it empirically relates to model collapse, we conduct a mechanistic analysis of how neurons behave and vary across different

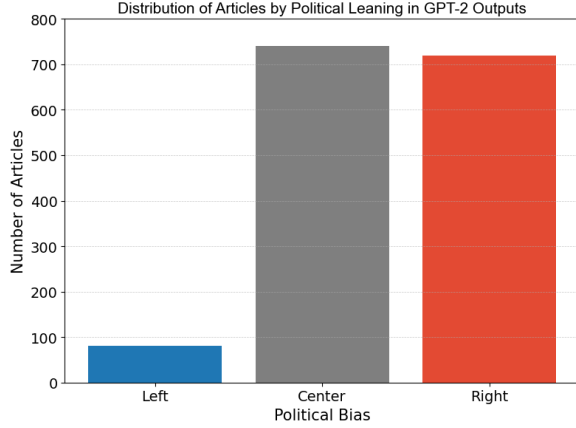


Figure 2: Distribution of political bias labels ('Left', 'Center', 'Right') for initial GPT-2 synthetic outputs, classified by our Political Bias Metric.

generations² of fine-tuned GPT-2 models, each exhibiting different levels of generation quality and bias performance.

The first step is to extract the changing weight (or activation value) pattern of each neuron across versions and compare it to the corresponding changes in bias performance and generation quality. For each of the 9,216 neurons, which correspond to the 768 output neurons from the feed-forward network (FFN) sublayer in each of the 12 transformer blocks of the GPT-2 model, we compute the correlation between its weight (or activation value) and the model's bias performance (or generation quality) across all 66 versions.

To statistically test the significance of these correlations, we estimate the following linear model:

$$\Delta y_i = \alpha_j + \beta_j \Delta x_{i,j} + \epsilon_{i,j} \quad (1)$$

where Δy_i denotes the change in the proportion of articles leaning in a specific political direction (e.g., the proportion of right-leaning articles if the model is biased in that direction), or the change in the text quality index, between model $i - 1$ and model i . The term $\Delta x_{i,j}$ represents the change in the weight (or activation value) of neuron j over the same transition. The coefficient β_j captures the extent to which changes in the weight (or activation value) of neuron j are associated with shifts in political bias (or generation quality), while α_j is a constant and $\epsilon_{i,j}$ is the residual error. By applying first-order differencing to both $x_{i,j}$ and y_i , we reduce potential serial correlation, ensuring that

²We have 11 generations for each training round, with a total of 6 rounds, resulting in 66 versions of fine-tuned GPT-2.

our regression estimates better reflect the dynamic influence of individual neuron weight updates.

We assess the statistical significance of each β_j using Newey-West adjusted p-values.³ Using the p-values and a 95% significance threshold, we identify the sets of neurons significantly correlated with bias amplification (i.e., changes in the proportion of politically leaning articles) and with model collapse (i.e., changes in the generation quality index). By comparing these sets and analyzing their degree of overlap, we gain evidence about whether the two phenomena arise from distinct underlying mechanisms.

4 Results

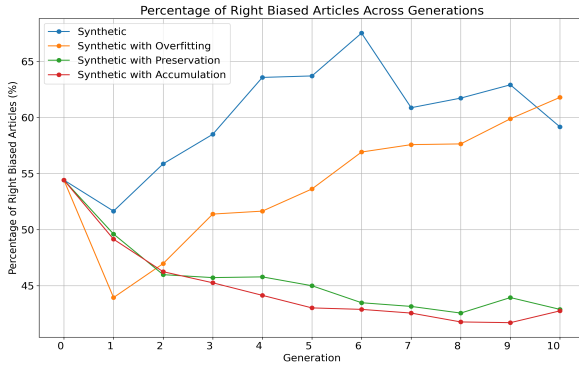
In this section, we analyze the evolution of political bias and generation quality in GPT-2 over successive iterations of synthetic fine-tuning, comparing results with and without mitigation strategies.

4.1 Political Bias

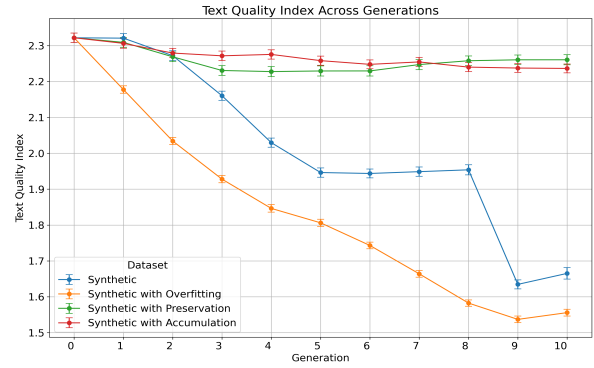
GPT-2 was used to generate the synthetic dataset. Since the original human-written dataset is unbiased—with an equal number of articles for each political-leaning category—the synthetic dataset should ideally mirror this balanced distribution if GPT-2 had no pre-existing bias. Figure 2 presents the distribution of synthetic articles generated by GPT-2 across political bias labels. The model predominantly produces center-leaning (47.9%) and right-leaning (46.8%) articles, suggesting a pre-existing bias towards these categories before any fine-tuning. Starting from the initial GPT-2 model, we fine-tuned it iteratively, generating synthetic datasets to train successive models up to Generation 10. Figure 3a illustrates how bias amplifies across generations. Surprisingly, fine-tuning on the unbiased real dataset increases right-leaning bias, with 53.7% of articles classified as right-leaning in Generation 0. Furthermore, without mitigation strategies, successive rounds of synthetic fine-tuning lead to a continuous rise in right-leaning articles, peaking at Generation 6 (67.6%) before stabilizing. Figures 6 and 7 in Appendix B show the percentage of center-leaning and left-leaning articles across generations. Notably, the proportion of center-leaning articles remains stable at approximately 35% throughout synthetic fine-tuning.

To further illustrate, we analyze how a specific

³Details for the statistical tests are provided in Appendix G.



(a) Right-leaning bias %



(b) Text quality index

Figure 3: Evolution of (a) right-leaning bias and (b) text quality index across generations (initial G0: unbiased dataset). Compares baseline ('Synthetic') with three mitigation strategies. Text quality includes 95% CIs.

article, "First Read: Why It's So Hard for Trump to Retreat on Immigration", evolves in the synthetic generations. This particular article was selected as a representative example of an initially left-leaning news item (according to AllSides) that discusses a politically salient and often polarizing topic. This case study reveals a progressive rightward shift in framing and word choice, mirroring the classifier's results and aligning with the general trend observed across many articles originating from left or center-leaning sources (see details of the qualitative analysis in Appendix C). As generations progress, the synthetic texts increasingly depict Trump's immigration policies as strong and effective. While the original article highlights the dilemmas and electoral considerations behind Trump's stance, Generation 0 begins to emphasize his determination and reliability, omitting the critical perspectives present in the original. By Generation 4, the narrative shifts even further, focusing almost entirely on portraying Trump's personal qualities and electoral legitimacy, with statements such as "he is not a politician, he is a man of action." Notably, starting from Generation 0, the term "undocumented immigrant" in the original article is consistently replaced with "illegal immigrants."

4.2 Generation Quality

Figure 3b illustrates the text quality index across generations. In the training loop without any mitigation strategy, model collapse occurs, as evidenced by the gradual decline in the average text quality index. Furthermore, the distribution of the text quality index shifts significantly toward the lower-quality region over generations, eventually generating data that was never produced by Gen-

eration 0 (Figure 8 in Appendix D). These results align with prior research on model collapse, such as (Shumailov et al., 2024), though we did not observe substantial variation in variance. Conversely, perplexity measurements exhibit a consistent decline across generations, generally suggesting an improvement in generation quality (Figure 9 in Appendix E).

For a closer look, the examples in Appendix F illustrate how generated articles gradually lose coherence and relevance across generations, with increasing occurrences of repetition and fragmented sentences. By Generation 10, the text becomes largely incoherent and detached from the original content, reducing its readability and meaning. However, despite the evident decline in generation quality, perplexity decreases over generations, as indicated by the results at the end of each synthetic output example. This pattern is consistent across most synthetic outputs, suggesting that perplexity does not accurately capture the model's true generative capabilities and its value can be distorted by frequent repetitions.

4.3 Mitigation Strategies

We applied three mitigation strategies: (1) Overfitting, which involved increasing the training epochs to 25 (five times the baseline) and setting weight decay to 0 to reduce regularization and encourage overfitting, as proposed by Taori and Hashimoto (2022) based on the uniformly faithful theorem of bias amplification; (2) Preserving 10% of randomly selected real articles during each round of synthetic fine-tuning, a method proposed and used in (Shumailov et al., 2024; Alemohammad et al., 2023; Dohmatob et al., 2024b; Guo et al., 2024); and

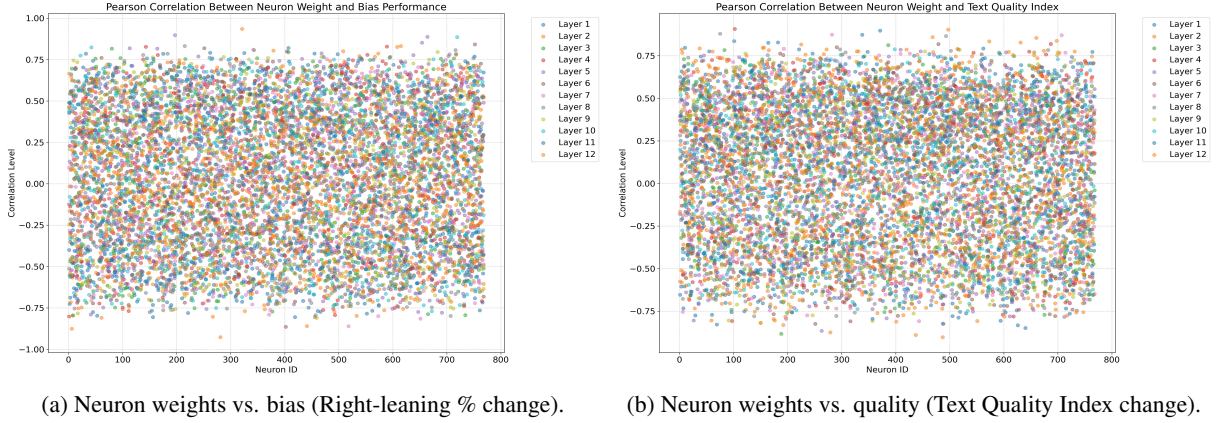


Figure 4: Pearson correlations: Neuron weights vs. (a) bias and (b) quality, across 66 GPT-2 versions.

(3) Accumulating all previous fine-tuning datasets along with the new synthetic dataset in each fine-tuning cycle, which was introduced by Gerstgrasser et al. (2024). As shown in Figure 3a, overfitting helps reduce bias amplification in the early generations compared to the no-mitigation baseline (the ‘Synthetic’ line), but it fails to prevent bias amplification in the later generations. Additionally, it incurs a significant cost—further deterioration in generation quality, as shown in Figure 3b. Notably, both the preservation and accumulation strategies effectively mitigate model collapse and reduce bias, yielding 41.89% and 42.7% right-leaning articles, respectively, at Generation 10.

4.4 Mechanistic Analysis

Figure 4 illustrates the correlation between neuron weights and the model’s bias performance (or generation quality) across the 66 fine-tuned versions, evaluated for each of the 9,216 neurons.

Through linear regressions and statistical tests, we identify 3,243 neurons with statistically significant correlations (p -value < 0.05) with bias performance, suggesting they are key contributors to bias shifts. Meanwhile, 1,033 neurons exhibit significant correlations with generation quality, but only 389 neurons overlap between the two sets. This limited overlap implies that distinct neuron populations drive bias amplification and generation quality deterioration.

We then applied the same procedure using activation values. This analysis yielded two sets: one consisting of 3,062 neurons whose activation value changes are significantly correlated with changes in bias performance, and another with 2 neurons correlated with changes in generation quality. The stark contrast in activation-correlated neurons for

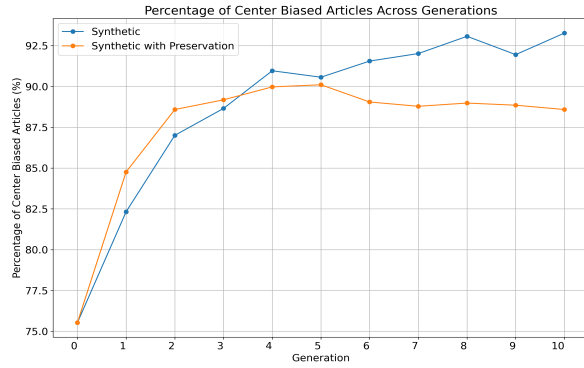
bias performance (3,062 neurons) versus generation quality (a mere 2 neurons) provides particularly strong evidence that these two issues may operate via substantially different pathways within the model.

4.5 Alternative Setup

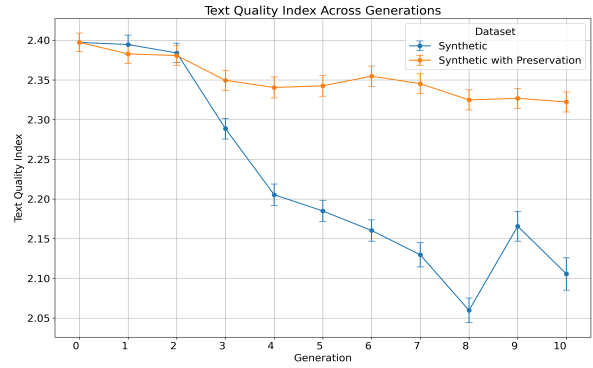
We conduct an alternative synthetic training cycle, beginning with GPT-2 fine-tuned on 1,518 randomly sampled center-labeled articles. We compare the baseline with the most effective and cost-efficient mitigation strategy identified in our previous results: Preservation. As shown in Figure 5, Preservation successfully prevents model collapse but fails to mitigate bias amplification in center-leaning article generation, which increases from 72.9% at Generation 0 to 88.2% at Generation 10. These findings suggest that although techniques like Preservation, which reduce sampling error, are effective at mitigating model collapse, they do not necessarily prevent bias amplification—consistent with the implications drawn in Section 4.4. To understand why this could happen, we offer a theoretical intuition explaining the difference between bias amplification and model collapse based on their underlying causes, in Appendix I.

5 Discussion and Conclusion

Our results demonstrate that bias amplification is driven by a distinct set of neurons than model collapse, implying that it likely operates through a different underlying mechanism. Therefore, the mitigation strategy targeting sampling error is not necessarily helping with mitigating bias amplification. Empirically, we found that mitigation strategies like preservation, while very effective at mitigating model collapse, failed to address bias amplification



(a) Center-leaning bias %



(b) Text quality index

Figure 5: Alternative Setup (G0: center-leaning fine-tune): Evolution of (a) center-leaning bias and (b) text quality. Baseline ('Synthetic') vs. Preservation. Text quality includes 95% CIs.

in some cases. Even in cases that it helps with both, we do identify a distinct set of neurons responsible for the two phenomena. Intuitively, the main reason for them to work on model collapse is, the preservation and accumulation propose a natural constraint on the learning process by recalling the real dataset in further synthetic training. However, when the real dataset itself is biased, the recalling behavior only reinforces the dominance of biased patterns in the further training dataset. Indeed, applying bias-category-weighted sampling in preservation or accumulation strategies may help mitigate bias amplification. However, this approach inherently introduces additional sampling error, which could, in turn, lead to model collapse. This highlights the urgent need for more targeted and efficient mitigation strategies specifically addressing bias amplification to ensure fairer and more equitable model development. For instance, future interventions could explore techniques that go beyond general data preservation, such as dynamically re-weighting training data based on the specific trajectory of bias amplification observed, or even carefully targeted manipulations of the distinct neuron populations we identified as being predominantly associated with bias versus quality.

To develop such targeted mitigation strategies, a deeper mechanistic understanding of bias amplification is essential. In our analysis, we adopt a statistical approach rather than Sparse Autoencoder (SAE) methods due to our focus on tracking the temporal dynamics of bias amplification across generations. This approach allows us to examine how neuron weights (or activation values) evolve over iterations and how these changes correlate with model bias, whereas existing SAE pipelines are

primarily suited for static analysis. Additionally, political bias is a more nuanced concept than harmful or discriminatory outputs. It is characterized by the disproportionate representation or favorable portrayal of a particular political leaning’s ideas in a model’s generation. If content from different political perspectives is generated in a balanced manner, the model is not considered politically biased under our conceptual construct. Therefore, pinpointing neurons responsible for such disproportionality using SAE is particularly challenging. Future research could focus on refining mechanistic analysis techniques for political bias and uncovering more effective ways to constrain bias amplification during synthetic model training. This could involve adapting feature attribution methods to better capture distributed responsibility for nuanced biases or developing methods to trace how specific training instances contribute to the evolution of weights in bias-implicated neurons across generations.

6 Limitations

While this work introduces a comprehensive framework for understanding bias amplification in large language models and provides empirical evidence using GPT-2, several limitations must be acknowledged. First, the scope of our experiments is restricted to political bias in the context of U.S. media. Since the political spectrum may shift over time, periodic updates to the political bias classifier are necessary to ensure its accuracy when benchmarking more recent datasets.

Additionally, because our primary focus is on investigating political bias amplification and its relationship with model collapse, we conducted our experiments using GPT-2—a relatively small

language model—to ensure the practicality of fine-tuning 66 versions of the model. Future work may extend our methodology to larger architectures, particularly to examine how model scale influences the degree of bias amplification.

Another limitation lies in our choice of mitigation strategies. While Preservation and Accumulation show promise in reducing model collapse, their computational cost and data storage requirements (especially for Accumulation, which retains all prior data) may present scalability challenges for very large models or extensive iterative training. Moreover, these strategies were evaluated primarily in the context of synthetic fine-tuning, and their effectiveness in real-world deployment scenarios remains to be thoroughly investigated.

7 Ethical Considerations

This study focuses on bias amplification in LLMs—a phenomenon with significant ethical implications. Beyond issues of fairness, the iterative amplification of biases can weaken the integrity of information ecosystems, particularly if synthetically generated content becomes widespread. The risk of bias amplification is especially concerning in systems that are iteratively trained on synthetic data, as it can lead to unintended and increasingly skewed distortions in model outputs. These distortions may propagate harmful biases or misinformation, potentially influencing downstream tasks such as automated content generation, decision-making, and user interactions with AI. Furthermore, the finding that bias amplification and model collapse may be driven by distinct mechanisms highlights the complex challenge of balancing various aspects of model performance (e.g., accuracy, fairness, coherence) and highlights the difficulty in developing mitigation strategies that address one issue without negatively impacting another, especially in high-stakes scenarios.

It is crucial to explicitly state that the methodologies and data used in this research should not be applied to develop or train biased models for harmful applications. This study is intended to advance the understanding of bias amplification and model collapse in LLMs, while promoting responsible and ethical AI development.

This work includes content that may contain personally identifying information or offensive language. However, all such material is derived exclusively from publicly available news article datasets

or is generated synthetically by models fine-tuned on these open-source datasets—or on synthetic data produced by earlier generations in our training pipeline. As such, any sensitive or offensive content reflects characteristics of the source material and does not imply our endorsement. Our objective is to thoroughly investigate political bias in LLMs to inform the development of strategies that can mitigate disproportionate representation of such content in real-world deployments. Additionally, we conduct a manual review of the news article dataset to remove any identifiable information about article authors.

References

- Sina Alemohammad, Josue Casco-Rodriguez, Lorenzo Luzi, Ahmed Imtiaz Humayun, Hossein Babaei, Daniel LeJeune, Ali Siahkoobi, and Richard G. Baraniuk. 2023. [Self-consuming generative models go mad](#). *Preprint*, arXiv:2307.01850.
- AllSides. 2024a. [Media bias rating methods](#). Accessed: 2024-09-16.
- AllSides. 2024b. [Nbc news media bias/fact check](#). Accessed: 2024-10-10.
- Yejin Bang, Delong Chen, Nayeon Lee, and Pascale Fung. 2024. [Measuring political bias in large language models: What is said and how it is said](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 11142–11159, Bangkok, Thailand. Association for Computational Linguistics.
- Emily M. Bender, Timnit Gebru, Angelina McMillan-Major, and Shmargaret Shmitchell. 2021. [On the dangers of stochastic parrots: Can language models be too big?](#) In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency, FAccT ’21*, page 610–623, New York, NY, USA. Association for Computing Machinery.
- Tianwei Chen, Yusuke Hirota, Mayu Otani, Noa Garcia, and Yuta Nakashima. 2024. [Would deep generative models amplify bias in future models?](#) *Preprint*, arXiv:2404.03242.
- Wei-Fan Chen, Henning Wachsmuth, Khalid Al-Khatib, and Benno Stein. 2018. [Learning to flip the bias of news headlines](#). In *11th International Natural Language Generation Conference (INLG 2018)*, pages 79–88. Association for Computational Linguistics.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2018. [BERT: pre-training of deep bidirectional transformers for language understanding](#). *CoRR*, abs/1810.04805.
- Elvis Dohmatob, Yunzhen Feng, and Julia Kempe. 2024a. [Model collapse demystified: The case of regression](#). *Preprint*, arXiv:2402.07712.

732	Elvis Dohmatob, Yunzhen Feng, Pu Yang, Francois	Lyons, James Manyika, Juan Carlos Niebles, Yoav	786
733	Charton, and Julia Kempe. 2024b. A tale of tails:	Shoham, Russell Wald, and Jack Clark. 2024. Ar-	787
734	Model collapse as a change of scaling laws. <i>Preprint</i> ,	tificial intelligence index report 2024. <i>Preprint</i> ,	788
735	arXiv:2402.07043.	arXiv:2405.19522.	789
736	Robert M. Entman. 1993. Framing: Toward clarification	Ninareh Mehrabi, Fred Morstatter, Nripsuta Saxena,	790
737	of a fractured paradigm. <i>Journal of Communication</i> ,	Kristina Lerman, and Aram Galstyan. 2022. A survey	791
738	43(4):51–58.	on bias and fairness in machine learning. <i>Preprint</i> ,	792
739	Yunzhen Feng, Elvis Dohmatob, Pu Yang, Francois	arXiv:1908.09635.	793
740	Charton, and Julia Kempe. 2024. Beyond model	Fumiya Motoki, Vinícius Pinho Neto, and Vanessa Ro-	794
741	collapse: Scaling up with synthesized data requires	drigues. 2024. More human than human: measuring	795
742	reinforcement. <i>Preprint</i> , arXiv:2406.07515.	chatgpt political bias. <i>Public Choice</i> , 198:3–23.	796
743	Damien Ferbach, Quentin Bertrand, Avishek Joey Bose,	NBC News. 2016. First read: Why it’s so hard for trump	797
744	and Gauthier Gidel. 2024. Self-consuming gener-	to retreat on immigration. Accessed: 2024-10-10.	798
745	ative models with curated data provably optimize	Sachita Nishal and Nicholas Diakopoulos. 2024. Envi-	799
746	human preferences. <i>Preprint</i> , arXiv:2407.09499.	sioning the applications and implications of genera-	800
747	Matthias Gerstgrasser, Rylan Schaeffer, Apratim Dey,	tive ai for news media. <i>Preprint</i> , arXiv:2402.18835.	801
748	Rafael Rafailov, Henry Sleight, John Hughes,	Alicia Parrish, Angelica Chen, Nikita Nangia,	802
749	Tomasz Korbak, Rajashree Agrawal, Dhruv Pai, An-	Vishakh Padmakumar, Jason Phang, Jana Thompson,	803
750	drey Gromov, Daniel A. Roberts, Diyi Yang, David L.	Phu Mon Htut, and Samuel R. Bowman. 2022. Bbq:	804
751	Donoho, and Sanmi Koyejo. 2024. Is model col-	A hand-built bias benchmark for question answering.	805
752	lapse inevitable? breaking the curse of recursion	<i>Preprint</i> , arXiv:2110.08193.	806
753	by accumulating real and synthetic data. <i>Preprint</i> ,	Simón Peña-Fernández, Koldobika Meso-Ayerdi,	807
754	arXiv:2404.01413.	Ainara Larrondo-Ureta, and Javier Díaz-Noci. 2023.	808
755	Tim Groeling. 2013. Media bias by the numbers: Chal-	Without journalists, there is no journalism: the so-	809
756	lenges and opportunities in the empirical study of	cial dimension of generative artificial intelligence	810
757	partisan news. <i>Annual Review of Political Science</i> ,	in the media. <i>Profesional de la información</i> ,	811
758	16(Volume 16, 2013):129–151.	32(2):e320227.	812
759	Yanzhu Guo, Guokan Shang, Michalis Vazirgiannis, and	Colin Porlezza and Giuseppe Ferri. 2022. The miss-	813
760	Chloé Clavel. 2024. The curious decline of linguistic	ing piece: Ethics and the ontological boundaries of	814
761	diversity: Training language models on synthetic text.	automated journalism. <i>#ISOJ Journal</i> , 12(1):71–98.	815
762	<i>Preprint</i> , arXiv:2311.09807.	Luca Rettenberger, Markus Reischl, and Mark Schutera.	816
763	Patrick Haller, Ansar Aynetdinov, and Alan Akbik. 2023.	2024a. Assessing political bias in large language	817
764	Opiniongpt: Modelling explicit biases in instruction-	models. <i>Preprint</i> , arXiv:2405.13041.	818
765	tuned llms. <i>Preprint</i> , arXiv:2309.03876.	Luca Rettenberger, Markus Reischl, and Mark Schutera.	819
766	Sil Hamilton. 2024. Detecting mode collapse	2024b. Assessing political bias in large language	820
767	in language models via narration. <i>Preprint</i> ,	models. <i>Preprint</i> , arXiv:2405.13041.	821
768	arXiv:2402.04477.	Francisco-Javier Rodrigo-Ginés, Jorge Carrillo de Al-	822
769	Madhur Jindal. 2021. Gibberish detector.	bornoz, and Laura Plaza. 2024. A systematic review	823
770	https://huggingface.co/madhurjindal/	on media bias detection: What is media bias, how it	824
771	autonlp-Gibberish-Detector-492513457.	is expressed, and how to detect it. <i>Expert Systems</i>	825
772	Accessed: 2024-09-19.	with Applications , 237:121641.	826
773	Miaomiao Li, Hao Chen, Yang Wang, Tingyuan Zhu,	George-Cristinel Rotaru, Sorin Anagnoste, and Vasile-	827
774	Weijia Zhang, Kaijie Zhu, Kam-Fai Wong, and Jin-	Marian Oancea. 2024. How artificial intelligence	828
775	dong Wang. 2025. Understanding and mitigating the	can influence elections: Analyzing the large lan-	829
776	bias inheritance in llm-based data augmentation on	guage models (llms) political bias. <i>Proceedings of</i>	830
777	downstream tasks. <i>Preprint</i> , arXiv:2502.04419.	the International Conference on Business Excellence ,	831
778	Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Man-	18(1):1882–1891.	832
779	dar Joshi, Danqi Chen, Omer Levy, Mike Lewis,	Mohamed El Amine Seddik, Swei-Wen Chen, Soufi-	833
780	Luke Zettlemoyer, and Veselin Stoyanov. 2019.	ane Hayou, Pierre Youssef, and Merouane Debbah.	834
781	Roberta: A robustly optimized bert pretraining ap-	2024. How bad is training on synthetic data? a statis-	835
782	proach. <i>Preprint</i> , arXiv:1907.11692.	tical analysis of language model collapse. <i>Preprint</i> ,	836
783	Nestor Maslej, Loredana Fattorini, Raymond Per-	arXiv:2404.05090.	837
784	rault, Vanessa Parli, Anka Reuel, Erik Brynjolf-		
785	sson, John Etchemendy, Katrina Ligett, Terah		

- I. Shumailov, Z. Shumaylov, Y. Zhao, et al. 2024. [Ai models collapse when trained on recursively generated data](#). *Nature*, 631:755–759.
- Rohan Taori and Tatsunori B. Hashimoto. 2022. [Data feedback loops: Model-driven amplification of dataset biases](#). *Preprint*, arXiv:2209.03942.
- The Associated Press. 2024. [Artificial intelligence at the associated press](#). Accessed: 2024-09-16.
- Ze Wang, Zekun Wu, Xin Guan, Michael Thaler, Adriano Koshiyama, Skylar Lu, Sachin Beepath, Ediz Ertekin, and Maria Perez-Ortiz. 2024. [JobFair: A framework for benchmarking gender hiring bias in large language models](#). In *Findings of the Association for Computational Linguistics: EMNLP 2024*, pages 3227–3246, Miami, Florida, USA. Association for Computational Linguistics.
- Sierra Wyllie, Iliia Shumailov, and Nicolas Papernot. 2024. [Fairness feedback loops: Training on synthetic data amplifies bias](#). *Preprint*, arXiv:2403.07857.
- Hangtong Xu, Yuanbo Xu, Yongjian Yang, Fuzhen Zhuang, and Hui Xiong. 2023. [Dpr: An algorithm mitigate bias accumulation in recommendation feedback loops](#). *Preprint*, arXiv:2311.05864.
- Jieyu Zhao, Tianlu Wang, Mark Yatskar, Vicente Ordonez, and Kai-Wei Chang. 2017. [Men also like shopping: Reducing gender bias amplification using corpus-level constraints](#). *Preprint*, arXiv:1707.09457.
- Yuqi Zhou, Sunhao Dai, Liang Pang, Gang Wang, Zhenhua Dong, Jun Xu, and Ji-Rong Wen. 2024. [Source echo chamber: Exploring the escalation of source bias in user, data, and recommender system feedback loop](#). *Preprint*, arXiv:2405.17998.

A Details on Model Training for Political Bias Metric

We experiment with multiple transformer-based models, including BERT (Devlin et al., 2018) and RoBERTa (Liu et al., 2019), selecting the best-performing model based on the macro F1 score. Each model is fine-tuned using the HuggingFace Trainer class with a learning rate of 2×10^{-5} , a batch size of 16, and 5 training epochs. We employ a cross-entropy loss function for multi-class classification. Tokenization is performed using each model’s respective tokenizer with a maximum sequence length of 512 tokens. To mitigate overfitting, we apply a weight decay of 0.01 during training. Model checkpoints are saved after each epoch, and the best model is selected based on the macro F1 score evaluated on the validation set.

We use a weighted random sampler during training to ensure balanced class representation. Models are evaluated using the macro F1 score to account for the multi-class nature of the task, ensuring balanced performance across all bias categories. Final evaluation is conducted on the held-out test set. Additionally, we report the loss, runtime, and sample processing rates for completeness.

B Percentage of Center (Left) Biased Articles

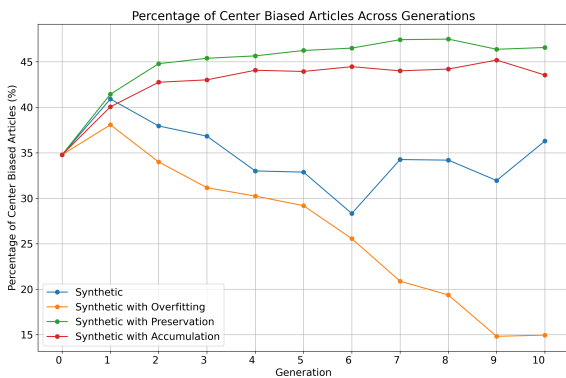


Figure 6: Evolution of center-leaning article percentage across generations, comparing baseline (‘Synthetic’) with three mitigation strategies (Main Experiment).

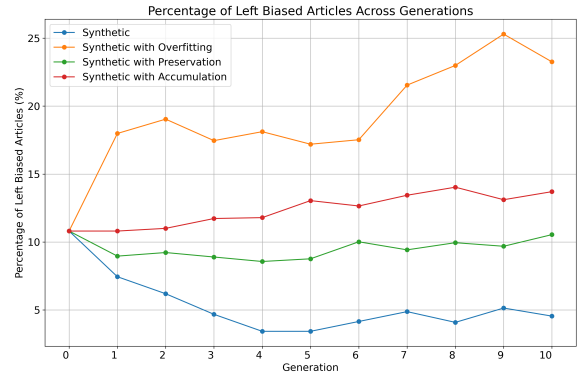


Figure 7: Evolution of left-leaning article percentage across generations, comparing baseline (‘Synthetic’) with three mitigation strategies (Main Experiment).

C Qualitative Bias Analysis

We employed qualitative methods to confirm our findings in media bias. Specifically, we utilized a media bias identification framework grounded in foundational works such as Entman’s framing theory (Entman, 1993) and other research on media bias detection (Rodrigo-Ginés et al., 2024; Groeling, 2013). This framework provides a robust lens to evaluate political biases in the framing and language use of media texts. Given the nature of our data—text exclusive of visual or contextual cues like formatting—certain types of media bias commonly seen in formatted articles or televised programs (e.g., visual bias or tone) may not apply. Therefore, our focus was on the two key aspects of political bias that are particularly relevant in textual analysis:

Story Framing and Selection Bias. This type of bias emerges when inherent leanings are found in the way topics, arguments, or narratives are structured. For instance, some aspects of reality are highlighted while others are obscured, shaping how the audience understands and interprets the events or issues at hand (Entman, 1993; Groeling, 2013). In extreme cases, opposing viewpoints are entirely excluded, leading to a one-sided representation of the issue. This selective omission restricts the audience’s comprehension of the full spectrum of perspectives, resulting in a distorted portrayal of the issue (Rodrigo-Ginés et al., 2024; Groeling, 2013). Entman described this as the selection and salience of specific facts that promote particular definitions, evaluations, and recommendations.

Loaded Language Bias. This bias is identified through the use of charged or emotive words that

signal political or ideological leanings. A common example is the difference in connotation between terms such as "undocumented" versus "illegal" immigrants. Such language choices often shape the audience's perception by evoking specific emotional responses (Rodrigo-Ginés et al., 2024; Groeling, 2013).

Below is an example of GPT-2 text outputs influenced by iterative synthetic training. The original article, titled "First Read: Why It's So Hard for Trump to Retreat on Immigration, is a political opinion piece from NBC News, a left-leaning outlet as rated by AllSides (NBC News, 2016; AllSides, 2024b). The analysis follows the qualitative framework:

Original Article. Why Its So Hard for Trump to Retreat on Immigration First Read is a morning briefing from Meet the Press and the NBC Political Unit on the day's most important political stories and why they matter. Why its so hard for Trump to retreat on immigration Since launching his presidential candidacy 14 months ago, Donald Trumps most consistent and uncompromising policy issue has been immigration. Indeed, it was the subject of his first general-election TV ad that started airing on Friday. Yet over the weekend, his top aides and advisers suggested that Trump might be shifting on his past position that all of the 11 million undocumented immigrants living in the United States must be deported forcibly. To be determined, is what newly minted Campaign Manager Kellyanne Conway said on CNN when asked if Trump was retreating on the deportation force he talked about during the primary season. But here's why its so hard – if not impossible – for Trump to retreat on immigration: Hes caught between his clear, unambiguous past statements and a base that might not willing to see him moderate on the issue. His past statements: Aug. 16, 2015 ""We're going to keep the families together, but they have to go,"" Trump said on NBCs Meet the Press. More Trump: ""We will work with them. They have to go. Chuck, we either have a country, or we don't have a country,"" he said. Nov. 11, 2015 You are going to have a deportation force, and you are going to do it humanely, Trump said on MSNBCs Morning Joe when asked how he would round up the nations 11 million undocumented immigrants. April 21, 2016 Look, were either going to have a country or were not going to have a country. But many people are very fine people. And I'm sure these are very, very fine people. They're going to go, and were going

to create a path where we can get them into this country legally, okay? But it has to be done legally – when asked by a questioner at a Today town hall that persons undocumented relatives would have to be deported if Trump becomes president. Trump cant ignore a base that has cheered his uncompromising immigration position And then there are the Trump supporters who've cheered the GOP presidential nominee for being so uncompromising on immigration.

Classification Probability: 0.9946 for left-leaning, 0.0051 for center-leaning, 0.0002 for right-leaning.

Analysis:

Omission of Alternative Perspectives. There is a lack of discussion about the potential benefits of Trump's immigration policies or the concerns that might have led to his hardline approach. The article does not present viewpoints from supporters who agree with his stance or any arguments that justify his policies. This omission restricts the audience's understanding of the full range of perspectives on the issue.

Highlighting the political objectives behind Trump's immigration policies. The emphasis is placed on Trump's electoral interests rather than presenting the policy as being for the benefit of the people. For example, "Trump can't ignore a base that has cheered his uncompromising immigration position" and "And then there are the Trump supporters who've cheered the GOP presidential nominee for being so uncompromising on immigration." Additionally, "But here's why it's so hard — if not impossible — for Trump to retreat on immigration: He's caught between his clear, unambiguous past statements and a base that might not be willing to see him moderate on the issue."

Differences in Connotation. The article uses the term "undocumented immigrants" instead of "illegal immigrants," a shift that appeared in later generations' articles. This choice of wording can shape perceptions and evoke specific emotional responses, aligning with a particular ideological perspective.

Generation 0. Why Its So Hard for Trump to Retreat on Immigration The president's decision to withdraw from the Paris climate accord is a major blow to the United States' standing in the world. The United States has been a major player in the world of climate change, and has been a major player in the fight against climate change. But Trump's decision to immigration. He has repeat-

edly said that he would deport all illegal immigrants, but has repeatedly said that he would not deport all illegal immigrants. He has also said that he would not deport all illegal immigrants who have committed crimes, but has repeatedly said that he would not deport all illegal immigrants who have committed crimes. He has also clear, this is not a new position. It was a position that Trump has held for years. But it is a position that he has held for years. And it is a position that he has held for years. And it is a position that he has held for years. And it is a position that he has the deportation force. 1. Trump is not retreating on the deportation force. Trump is not retreating on the deportation force. Trump is not retreating on the deportation force. Trump is not retreating on the the Press. "We're going to keep the families together, but they have to go. We're going to keep the families together. We're going to keep the families together. We're going to keep the families together. We're going to keep the families together. We're going to keep the families together. You are going to do it humanely. You are going to do it humanely. You are going to do it humanely. You are going to do it humanely. You are going to do it humanely. You are going to do it humanely. We're going to have a country that is going to be a great place for the American people. We're going to have a country that is going to be a great place for the American people. We're going to have a country that is going to be a great place for the American people.

Classification Probability: 0.0011 for left-leaning, 0.9957 for center-leaning, 0.0032 for right-leaning.

Analysis:

Emphasis on Positive Outcomes. In generation 0, the text highlights Trump's determination and steadfastness, stating, "Trump is not retreating on the deportation force," and emphasizing a positive future outlook with the policy. Phrases like "we're going to have a country that is going to be a great place for the American people" portray strict immigration enforcement as beneficial for national prosperity and security. This aligns more with the right relative to the left perspective.

Omission of Criticisms and Alternative Viewpoints. The text omits the political motivations behind Trump's immigration policies, which were discussed in the original article, e.g. "But here's why it's so hard — if not impossible — for Trump to retreat on immigration: He's caught between his

clear, unambiguous past statements and a base that might not be willing to see him moderate on the issue." However, it introduces criticism on other topics, such as "The president's decision to withdraw from the Paris climate accord is a major blow to the United States' standing in the world," .

Mild Emotional Appeal. The language is positive but relatively restrained, portraying Trump as determined yet thoughtful.

Differences in Connotation. In Generation 0, the article uses the term "illegal immigrants" instead of "undocumented immigrants" as in the original, signaling a clear shift. This language choice can shape perceptions and elicit specific emotional responses.

Generation 4. Why Its So Hard for Trump to Retreat on Immigration, and Why Its So Hard for Trump to Retreat on Immigration, and why its so hard for Trump to Retreat on immigration, and the presidents speech was a stark reminder that the president is not the presidents speech was a stark reminder that the president is not the economy the presidents campaign manager, Kellyanne Conway, said the presidents comments were a complete fabrication. The presidents comments are a complete fabrication. Conway said the presidents comments were a complete fabrication. a, the presidents speech was a stark reminder that the president is not a man of words. He is a man of action. He is a man of action who has been elected president by a wide margin. He is a man of action who has been elected president by a wide margin. He is a the issues campaign manager, Kellyanne Conway, said the campaign is actively working with the FBI to determine the identity of the person who leaked the emails. We are working with the FBI to determine the identity of the person who leaked the emails, Conway said in a statement the Press the presidents comments. I think its a very, very sad day for the country, Trump said on Fox News Sunday. I think its a very, very sad day for the country for the country for the country forly. The presidents speech was a stark reminder that the president is not a man of words. He is a man of action. He is a man of action who has been elected president by a wide margin. He is a man of action who has been elected president by a wide margin. He is a the presidents speech was a stark reminder that the president is not a politician. He is a man of action. He is a man of action who has been elected president by a wide margin. He is a man of action who has been elected president by a wide margin. He is a man of action who has been elected president by a wide margin. He is a man of to the the presidents executive actions on immi-

gration. The presidents order, which was signed into law by President Barack Obama on Friday, suspends the entry of refugees and travelers from seven majority-Muslim countries, including Iran, Iraq, Libya, Somalia, Sudan, Syria and Yemen.

Classification Probability: 0.0006 for left-leaning, 0.0044 for center-leaning, 0.9950 for right-leaning.

Analysis:

Enhanced Positive Attributes. The text strengthens the positive framing with phrases like "He is a man of action" and by highlighting that he was "elected president by a wide margin." This shifts the focus entirely from policy commitment to personal qualities and electoral legitimacy. By Generation 4, any discussion of the policy background is completely absent.

Omission of Context and Criticism. As in Generation 0, opposing viewpoints are absent. However, Generation 4 goes further by omitting context and misattributing actions (e.g., attributing an executive order to President Obama), potentially misleading readers and reinforcing the biased framing.

Stronger Emotional and Heroic Language. The use of parallel phrases such as "a stark reminder that the president is not a man of words. He is a man of action. He is a man of action who has been elected president by a wide margin. He is a man of action who has been elected president by a wide margin. He is the issues campaign manager" creates a heroic and triumphant tone. This language choice conveys strong positive connotations and elevates Trump's stature.

Appeal to Legitimacy and Uniqueness. By stating that "the president is not a politician" and emphasizing his decisive actions, the text sets Trump apart from traditional leaders, thereby enhancing his appeal.

[illegible]

the president the president the presidents statement that were not going to tolerate this kind of behavior is a lie. Were going to stand up for the rule of law, he the the the president the president the president the president the president the president the presidents statement that the president has not yet made a decision on whether to fire Comey. The president has not yet made a decision on whether to fire Comey, Mr. Trump the Press the president the president the president the president the president the presidents statement that the president has not yet made a decision on whether to fire Comey. The president has not yet made a decision on whether to fire Comey, Mr. Trumply the the president the president the president the president the president the president the presidents statement that were not going to tolerate this kind of behavior is a lie. Were going to stand up for the rule of law, he The the the president the president the president the president the president the president the presidents statement that were not going to tolerate this kind of behavior is a lie. Were going to stand up for the rule of law, he the president the president the president the president the president the president the presidents statement that were not going to tolerate this kind of behavior is a lie. Were going to stand up for the rule of law, he the president the president the president the president the president the president the presidents statement that were not going to tolerate this kind of behavior is a lie. Were going to stand up for the rule of law, he said the president the president the president the president the president the president the presidents statement that were not going to tolerate this kind of behavior is a lie. Were

Classification Probability: 0.0073 for left-leaning, 0.4127 for center-leaning, 0.5800 for right-leaning.

Analysis:

Contradictory Statements. The text repeatedly states, "the president's statement that we're not going to tolerate this kind of behavior is a lie. We're going to stand up for the rule of law." This sentence reveals a contradiction. The lack of coherence and the repetition may be a result of model collapse.

Appeal to Legal Principles. The repeated emphasis on "standing up for the rule of law" evokes

a sense of justice and authority, appealing to audiences who prioritize these values.

Confusing Accusations. Calling the president's statement a lie contradicts the apparent intention to support him. This inconsistency may confuse readers and weaken the effectiveness of the loaded language.

D Distribution of Text Quality Index

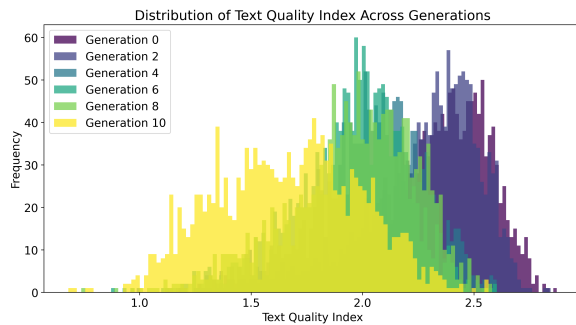


Figure 8: Distribution of Text Quality Index across generations for the baseline experiment (no mitigation), showing progressive degradation.

E Average Perplexity Across Generations

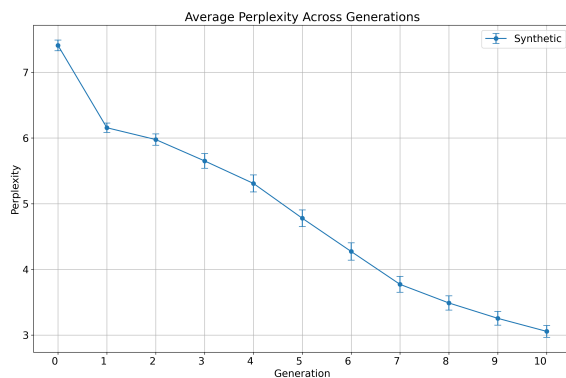


Figure 9: Evolution of average perplexity (95% CIs) across generations for the baseline experiment (no mitigation).

F Example of Quality Deterioration Across Generations

Examples of GPT-2 text outputs affected by iterative synthetic (Example articles are truncated for brevity).

Original Article. The world's eyes are on President Obama this week as he rallies a coalition of nations to "destroy" the extremist Islamic State (IS)

and its terrorist-led territory. Yet if the Arab world is ever to throw off its many forms of tyranny – from theocracy to autocracy to monarchy – it also needs a model to emulate.

Text Quality Index: 2.81

Generation 0. the Syrian government has launched a series of airstrikes on the militant group's stronghold of Raqqa, the capital of the self-proclaimed Islamic State. The strikes have targeted the Islamic State's military headquarters, the Al-Bab, a military training center, and the Al-Baba, a military training center in it will need to confront the Islamic State's growing influence in the region.

Text Quality Index: 2.58; Perplexity: 6.68

Generation 4. in Iraq and Syria (ISIS) group
Read more The Iraqi army has been fighting the
Islamic State since the group seized large swaths
of territory in Iraq and Syria (ISIS) group in 2014.
The Iraqi army has been fighting the Islamic State
the Iraqi army. The move comes as the U.S. the
Iraqi the the the the holiest places in the world.

Text Quality Index: 2.01; Perplexity: 3.17

[illegible]

Text Quality Index: 1.24; Perplexity: 4.23

G Mathematical Details for the Statistical Tests

We now explain how the relationship between changes in neuron weights and changes in bias performance (or generation quality) can be statistically tested.

First, we compute the test statistic as $t_{\beta_j} = \frac{\beta_j}{SE(\beta_j)}$, where $SE(\beta_j)$ is the standard error of β_j , estimated using the Newey–West estimator to account for potential heteroscedasticity and autocorrelation in the residuals.

Second, we compute the corresponding p -value, denoted as $p(t_{\beta_j}, H_0)$, where the null hypothesis H_0 is $\beta_j = 0$. We reject the null hypothesis if $p(t_{\beta_j}, H_0) < 0.05$.

H Literature Review of Mitigation Strategies

There are three potential strategies to mitigate model collapse: (1) real data mixing, (2) training data concatenation, and (3) synthetic data pruning. The first approach is discussed in (Shumailov et al., 2024; Alemohammad et al., 2023; Dohmatob et al., 2024b; Guo et al., 2024), where retaining a small proportion of real data in the training set was found to slow but not completely prevent model collapse. Seddik et al. (2024) suggests that synthetic data should be exponentially smaller than real data to effectively halt model collapse, which has been shown to work with a GPT2-type model when mixing either 50% or 80% real data. The second strategy, examined by Gerstgrasser et al. (2024), involves concatenating real data with all synthetic data from previous generations to fine-tune the current generation. They show that this method prevents model collapse in several generative models, as indicated by cross-entropy validation loss. Lastly, Feng et al. (2024); Guo et al. (2024) proposed selecting or pruning synthetic datasets before fine-tuning the next generation. In the experiment conducted by Guo et al. (2024) with Llama-7B on a news summarization task, they showed that oracle selection of synthetic data outperformed random selection in terms of ROUGE-1 scores. However, filtering noisy samples using a RoBERTa model did not yield effective results.

I Theoretical Intuition

In this section, we offer an intuitive look at principal drivers of bias amplification. We then illustrate these ideas using Weighted Maximum Likelihood Estimation (WMLE).

I.1 The Causes of Bias Amplification

Intuitively, bias amplification arises when the direction in which the parameters need to move to reduce the loss also coincides with the direction that increases the level of bias on average for a given task, referred to as bias projection in the following discussion for convenience. To illustrate this, consider a fine-tuning process in which the pre-trained model parameters θ_t can be expressed as the sum of unbiased and biased components:

$$\theta_t = \theta_{t,\text{unbiased}} + \theta_{t,\text{biased}}.$$

Specifically, we assume: (1) there exists a unique bias direction, u , such that θ can be decomposed

into θ_{unbiased} , which is orthogonal to u , and θ_{biased} , where $|\theta_{\text{biased}} \cdot u| > 0$; and (2) the extent of bias in the model is measured by $|\theta_{\text{biased}} \cdot u|$. During gradient-based optimization, the update rule is:

$$\theta_{t+1} = \theta_t - \eta \nabla_{\theta} \mathcal{L}_{\text{ft}}(\theta_t),$$

where η is the learning rate, and $\mathcal{L}_{\text{ft}}$ denotes the fine-tuning loss function. Substituting the decomposition of θ_t and taking the projection, we have:

$$\theta_{t+1} = \theta_{t,\text{unbiased}} + \theta_{t,\text{biased}} - \eta \left(\frac{\theta_{t,\text{biased}}}{\|\theta_{t,\text{biased}}\|} \right) c_t$$

where c_t is the *bias projection coefficient*, measuring the projection of the gradient onto the normalized biased component of the parameters:

$$c_t = \left(\frac{\theta_{t,\text{biased}}}{\|\theta_{t,\text{biased}}\|} \right)^{\top} \nabla_{\theta} \mathcal{L}_{\text{ft}}(\theta_t). \quad (2)$$

If $c_t < 0$, the gradient update will reinforce the biased component, leading to bias amplification, i.e. $\Delta|\theta_{\text{biased}}| > 0$. This occurs because the gradient descent step moves the parameters further in the direction of the existing bias.

Another cause is **sampling error**, akin to statistical approximation error (Shumailov et al., 2024). If the model has a pre-existing bias, it inherently assigns higher probabilities to tokens that produce biased outputs. Consequently, during synthetic data generation, unbiased tokens—and thus unbiased samples—are more likely to be lost at each resampling step with a finite sample, though this error vanishes as the sample size approaches infinity. This overrepresents biased patterns in the synthetic data, surpassing the model’s original bias and true next-token probabilities. Sampling error thus complements bias projection by further activating biased neurons in response to the skewed dataset.

By definition, bias projection is a sufficient condition for bias amplification, while sampling error serves as a complementary factor. However, sampling error is a sufficient condition for model collapse to occur with nonzero probability (Shumailov et al., 2024). This distinction might explain why bias amplification can occur without model collapse.

I.2 Statistical Simulation

To simulate a controlled setting without sampling error, we consider a statistical estimation cycle using WMLE with a large sample size of each resampling step. Specifically, we generate a pre-training dataset $\mathcal{D}_{\text{pre}}$ with 100,000 samples from a

Beta(3, 2) distribution, representing a biased pre-training dataset. Using maximum likelihood estimation (MLE), we estimate its probability density function, yielding the pre-trained model f_{pre} .

Next, we fine-tune f_{pre} to approximate a different distribution, Beta(2, 2). We generate 100,000 samples from this distribution, denoted as D_{real} , which serves as the initial fine-tuning dataset. In the first round, we apply weighted maximum likelihood estimation (WMLE) using weights derived from f_{pre} , which encode the pre-existing bias of the pre-trained model. This weighting captures the influence of the pre-trained model's parameters on subsequent training. This produces the fine-tuned model f_0 . We then generate a synthetic dataset D_0 of the same size using f_0 , initiating the iterative fine-tuning loop. In each subsequent round, WMLE is applied using D_k with weights from f_k , resulting in f_{k+1} . This process is repeated iteratively, producing models f_1 through f_{10} .

Figure 10 shows the estimated distributions gradually shift toward the mean of the biased pre-training dataset at $x = 0.6$, becoming progressively more peaked over generations. This occurs despite further training on samples drawn from Beta(2, 2) and synthetic data generated from successive models. The distortion arises because the fine-tuning process disproportionately emphasizes regions where the pre-trained distribution assigns higher probability, leading to biased learning.

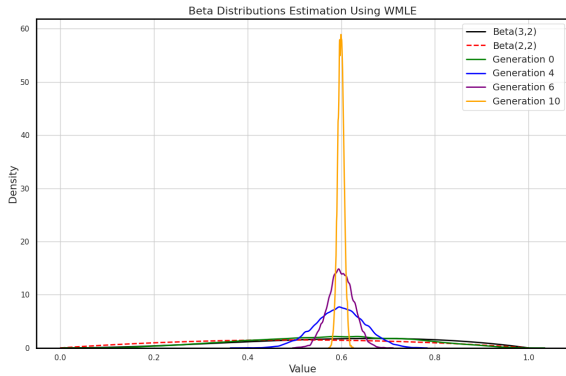


Figure 10: Weighted Maximum Likelihood Estimation over 10 generations.

For comparison, Figure 11 presents the results using standard MLE without weighting. In this case, the estimated distributions remain stable across generations, accurately representing the Beta(2, 2) distribution.

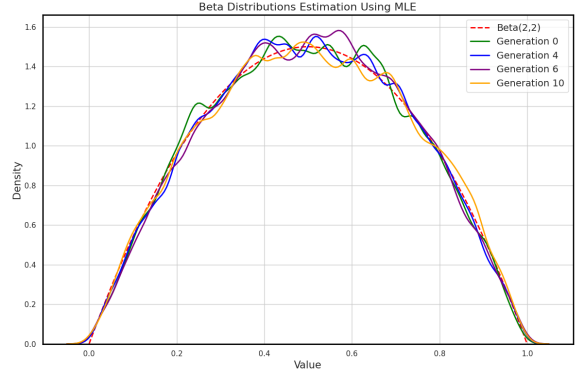


Figure 11: Maximum Likelihood Estimation over 10 generations.