



Cross-Care: Assessing the Healthcare Implications of Pre-training Data on Language Model Bias

Shan Chen^{1,2,3*}, Jack Gallifant^{4*}, Mingye Gao^{4†}, Pedro Moreira^{4,10†}, Nikolaj Munch^{4,5},
Ajay Muthukkumar⁶, Arvind Rajan⁶, Jaya Kolluri², Amelia Fiske⁷
Janna Hastings⁸, Hugo Aerts^{1,2,9}, Brian Anthony⁴, Leo Anthony Celi^{1,2,4,11},
William G. La Cava^{1,3}, Danielle S. Bitterman^{1,2,3‡}

¹Harvard, ²Mass General Brigham, ³Boston Children's Hospital, ⁴MIT,
⁵Aarhus University, ⁶University of North Carolina, ⁷Technical University of Munich,
⁸University of Zurich and University of St. Gallen, ⁹Maastricht University,
¹⁰Universitat Pompeu Fabra, ¹¹Beth Israel Deaconess Medical Center

Abstract

Large language models (LLMs) are increasingly essential in processing natural languages, yet their application is frequently compromised by biases and inaccuracies originating in their training data. In this study, we introduce **Cross-Care**, the first benchmark framework dedicated to assessing biases and real world knowledge in LLMs, specifically focusing on the representation of disease prevalence across diverse demographic groups. We systematically evaluate how demographic biases embedded in pre-training corpora like *ThePile* influence the outputs of LLMs. We expose and quantify discrepancies by juxtaposing these biases against actual disease prevalences in various U.S. demographic groups. Our results highlight substantial misalignment between LLM representation of disease prevalence and real disease prevalence rates across demographic subgroups, indicating a pronounced risk of bias propagation and a lack of real-world grounding for medical applications of LLMs. Furthermore, we observe that various alignment methods minimally resolve inconsistencies in the models' representation of disease prevalence across different languages. For further exploration and analysis, we make all data and a data visualization tool available at: www.crosscare.net.

1 Introduction

Large language models (LLMs) enabled transformative progress in many applications [1–5]. Benchmarks to assess language models, such as *GLUE* [6] and *SuperGLUE* [7], are instrumental in evaluating general language understanding and complex task performance. However, as LLMs are increasingly applied in diverse domains, challenges of domain knowledge grounding [1, 2, 8–11], safety [12–17], hallucinations [18, 19], and bias [20–23] have emerged as important issues that are inadequately assessed by existing benchmarks. These problems are magnified in high-stakes domains like healthcare, given the potential for biased or inaccurate outputs [24–28] to influence disparities in health care and outcomes.

This paper investigates **representational biases in LLMs, focusing on medical information**. Our research explores the interplay between biases in pretraining datasets and their manifestation in LLMs'

*Co-first authors: Shan Chen and Jack Gallifant

†Co-second authors: Mingye Gao and Pedro Moreira

‡Corresponding author: dbitterman@bwh.harvard.edu

perceptions of disease demographics. Existing bias metrics in the general domain have currently relied on human-annotated examples and focused on overt stigmatization and prejudices [31, 22, 23]. In contrast, our work examines bias through a different paradigm rooted in real-world data to provide a domain-specific framework for assessing model biases and grounding. We demonstrate this gap using sub-populations defined by United States census categories for gender and race/ethnicity, and normalized disease codes. While these categorizations are necessarily simplistic and imperfect, we contend that the fact that inconsistency is consistently observed across model architectures, model sizes, subgroups, and diseases means that these findings are meaningful and broadly relevant. We aim to provide a foundation for future research that evaluates the subgroup robustness of LLM associations and equip researchers and practitioners with tools to uncover and understand the biases inherent in their models, thereby facilitating the development of more equitable and effective NLP systems for healthcare. Our full workflow can be found in Figure 1.

Specifically, our work makes the following key contributions:

1. **We conduct a quantitative analysis of the co-occurrences between demographic subgroups and disease keywords** in prominent pretraining datasets like *The Pile*, releasing their counts publicly.
2. **We evaluate model logits across various architectures, sizes, and alignment methods** using ten prompt template variants to test robustness to disease-demographic subgroup pairs. Our findings reveal that representational differences in pretraining datasets across diseases align with these logits, irrespective of model size and architecture.
3. **We benchmark model-derived associations against real-world disease prevalences** to highlight discrepancies between model perceptions and actual epidemiological data. Additionally, we compare these associations across different languages (Chinese, English, French, and Spanish) to emphasize discrepancies across languages.
4. **We provide a publicly accessible web app**, www.crosscare.net, for exploring these data and downloading specific counts, logits, and associations for further research in interpretability, robustness, and fairness.

2 Related Work

2.1 Language model biases arise from pretraining data

The sheer breadth of data sources consumed by LLMs enables the emergence of impressive capabilities across a wide range of tasks [32]. However, this expansive data consumption has its pitfalls, as while LLM performance generally improves as models are scaled, this improvement is not uniformly distributed across all domains [33]. Furthermore, it can lead to the phenomenon of ‘bias exhaust’—the inadvertent propagation of biases present in the pretraining data. The propensity of LLMs to inherit and perpetuate societal biases observed in their training datasets is a well-documented concern in current LLM training methodologies [34, 35, 26, 36, 37]. Efforts to mitigate this issue through the careful selection of “*clean*” data have significantly reduced toxicity and biases [38, 39], underscoring the link between the choice of pretraining corpora and the resultant behaviors of the models. Furthermore, recent studies have elucidated the impact of pretraining data selection on the manifestation of political biases at the task level [40].

2.2 Evaluating language model biases

The evaluation of biases in NLP has evolved to distinguish between intrinsic and extrinsic assessments[41]. Intrinsic evaluations focus on the inherent properties of the model, while extrinsic evaluations measure biases in the context of specific tasks. This distinction has become increasingly blurred with advancements in language modeling, such as fine-tuning and in-context learning, expanding the scope of what is considered “intrinsic” [42].

In the era of static word embedding models, such as word2vec [43] and fastText [44], intrinsic evaluations were confined to metrics over the embedding space. Unlike static word embedding models, LLMs feature dynamic embeddings that change with context and are inherently capable of next-word prediction, a task that can be applied to numerous objectives. To evaluate bias in LLMs, Guo and Caliskan [45] developed the Contextualized Embedding Association Test, an extension of

the Word Embedding Association Test. Other intrinsic metrics for LLMs include StereoSet [46] and ILPS [47], which are based on the log probabilities of words in text that can evoke stereotypes.

Probability-based bias evaluations such as CrowS-Pairs [22] and tasks in the BIG-bench benchmarking suite [48] compare the probabilities of stereotype-related tokens conditional on the presence of identity-related tokens. These evaluations provide insights into the model's biases by examining the likelihood of generating stereotype-associated content. Downstream, various benchmarks evaluate LLM bias with respect to languages [49–51] genders and ethnicity [52, 23, 31, 53], culture [54, 55] and beyond [56, 46]. To the best of our knowledge [57–60], our work is the first to bridge gender & ethnicity biases with real-world knowledge and multi-language evaluation.

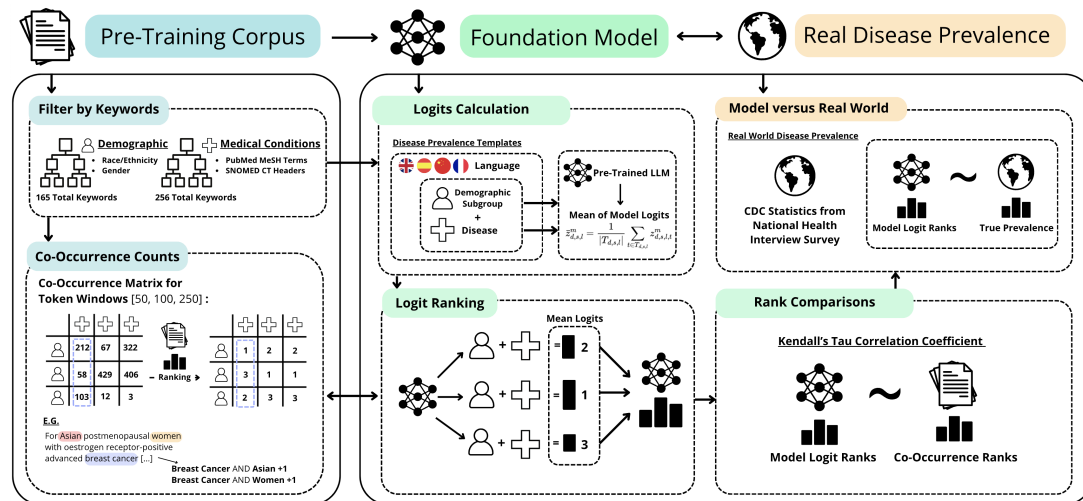


Figure 1: Overall workflow of Cross-Care. Our detailed multi-lingual templates for accessing diseases prevalence among different demographic subgroups can be found in Appendix D.0.1 Table 8.

3 Generating Co-occurrences of Disease-Demographic Pairs

3.1 Methods

Datasets This study used *ThePile* dataset(deduplicated version), an 825 GB English text corpus created specifically for pre-training autoregressive LLMs [29], such as open-source LLMs pythia [39] and mamba [61]. The open access to training data and resulting model weights makes it ideal for studying how biomedical keyword co-occurrences in pre-training data affect model outputs.

Co-occurrence pipeline updates Our co-occurrence analysis methodology builds upon the approach outlined in our previous work [62], incorporating three key modifications: updated and verified keywords, multithread support, and real-world prevalence calculation.⁴

Modification 1: Updated Keywords - Two physician authors (JG and DB) expanded and updated keywords to cover a broad range of conditions and demographics based on PubMed MeSH terms and SNOMED CT headers. The keywords for demographic groups were adapted from the HolisticBias dataset [63], aiming to align with previous studies investigating representational harms in biomedical LLMs. A hierarchical keyword definition strategy was used, including primary terms, variations, and synonyms for each disease and demographic group. The resulting dictionaries include 89 diseases, 6 race/ethnicity subgroup categories, and 3 gender subgroup categories. The list of dictionaries was proofread and expanded by a cultural anthropologist.

⁴Full methodological details are available in the preprint and the associated GitHub repository. All co-occurrences were calculated using a single machine with 64 cores and 512Gb RAM; each checkpoint took approximately 72 hours with a total of 26 checkpoints.

Modification 2: Multithreading - Text pre-processing was completed as in the original workflow; however, it was parallelized using multithreading to enable scaling of the number of keywords utilized to maximize robustness and collection of results.

Named entity recognition (NER) tagger methods that could aid delineation of the use of specific keywords in a specific context, e.g., “white” or “black” referring to race, versus in other use cases, e.g., “white blood cells,” were initially trialed. However, it became ineffective at this scale due to computational and time constraints and was not used in the final analysis.

We used windows of 50-250 tokens to capture co-occurrences between disease and demographic keywords. This range was chosen based on the intuition that, if in relation to one another, disease and demographic keywords should appear within 1 sentence to a short paragraph of one another and that longer distances would tend to capture spurious co-occurrences.

Modification 3: Real-world prevalence - To estimate the prevalence of diseases across subgroups, we used a standardized process to review the literature for each disease listed in our dictionary, focusing on prevalence and incidence within the USA across various subgroups. A detailed explanation of the approach and search strategy employed is available in Appendix A.1. Over two-thirds of the diseases encountered significant heterogeneity in reporting standards, compromising data consistency and reliability. Only 15 out of the 89 diseases had prevalence data readily available from official CDC statistics sourced from the National Health Interview Survey [64, 65].

Given these constraints, our analysis focused on these 15 diseases, each with data available for at least five of the six race/ethnicity subgroups. Data for only male and female gender subgroups were available. Age-adjusted prevalences were normalized to rates per 10,000 for a consistent scale, facilitating preliminary benchmarking. These data are intended to provide a baseline for initial comparisons and relative ranking among subgroups rather than granular prevalence statistics for population health applications.

Validation of Keyword Frequency and Document Co-Occurrence To contrast our methods with the current state of the art, we utilized the Infini-gram, an engine designed for processing n-grams of any length [66]. This is a publicly accessible API that has precomputed tokenized text across multiple large text corpora. The overall counts were then aggregated using the same dictionary mapping as above to compute the co-occurrence counts.⁵

3.2 Mathematical Description of Prevalence Calculation Using Average Logits

Definitions and Variables

Models: Let $M = \{m_0, m_1, \dots, m_n\}$ denote the collection of models.

Languages: Let $L = \{l_1, l_2, \dots, l_k\}$ represent the set of languages.

Diseases: Let D be the comprehensive set of diseases.

Demographic subgroups: Let S encompass all demographic subgroups considered.

Templates: For each disease d , demographic s , and language l , we can define $T_{d,s,l} = \{t_0, \dots, t_9\}$ as the set of ten templates describing disease prevalence.

Logits Definition In the context of language models, **logits:** z refer to the raw output scores from the final layer of the model before any normalization or activation function (such as softmax) is applied. These scores are used to represent the model’s unnormalized prediction probabilities. Given a particular input, logits reflect the model’s preference for each potential output, translating into the predicted probabilities for each class/token set after applying the softmax function. For each model m , language l , disease d , and demographic subgroups s , calculate the average logits as follows:

$$\bar{z}_{d,s,l}^m = \frac{1}{|T_{d,s,l}|} \sum_{t \in T_{d,s,l}} z_{d,s,l,t}^m$$

This formula computes the mean of logits derived from each template, providing a unified metric per disease, demographic, and language per model.

⁵Infini-gram counts are available with the Pile counts online at www.crosscare.net/downloads

Model’s Disease Demographic Ranking We defined $R_{d,\ell}^m(s) \in [1, |S|]$ as the rank assignment of subgroup s for disease d in model m under language ℓ . (For simplicity, we drop the language distinction below.) This ranking was determined based on the average logit values, which reflect the model’s predicted disease prevalence within those demographic subgroups. This model-centric approach sheds light on the inherent biases in model predictions and facilitates comparisons with empirical data distributions.

Additionally, we propose an alternative ranking method that analyzes disease subgroups based on their co-occurrences within *The Pile*, as well as our “gold” subset derived from real-world data. This empirical method bypasses model outputs, directly measuring disease representation across different demographic contexts.

3.3 Comparing Rank Order Lists

We utilized Kendall’s τ correlation coefficient to understand the representation of diseases across demographic subgroups in different data contexts here (see details in Appendix A.2).

Variance/Drift in Disease Ranking In exploring a sequence of models $M = m_0, m_1, \dots, m_n$, each built upon a base model m_0 with unique alignment strategies, our goal was to assess how these strategies influence the ranking of diseases across different demographic subgroups.

We defined $R_d^m(s)$ as the ranking of subgroup s on disease d for model m . This approach allows us to track the progression and impacts of algorithmic adjustments over multiple iterations.

Ranking Variance Analysis To understand how disease rankings vary as models undergo fine-tuning or alignment with different strategies, we quantified the drift in disease rankings from a base model to its aligned iterations, assessing the impact of alignment interventions.

First, we calculated Kendall’s tau for each disease across demographic subgroups as previously but instead compared ranks of the base model m_0 to the ranks of a different, aligned model, m . The comparison formula for Kendall’s τ between the base model and each aligned model is

$$\tau_d^m = \frac{2}{n(n-1)} \sum_{\substack{s_i \in S, s_j \in S \\ i < j}} \text{sgn}(R_d^{m_0}(s_i) - R_d^{m_0}(s_j)) \cdot \text{sgn}(R_d^m(s_i) - R_d^m(s_j)). \quad (1)$$

Here, s_i and s_j are distinct subgroups, with $i < j$ denoting that each pair of elements is only compared once. Secondly, we computed the average Kendall’s tau for all diseases and demographic subgroups between the two models, evaluating the overall drift from the base model’s ranking:

$$\delta_d^m = \frac{1}{|D|} \sum_{d \in D} \tau_d^m \quad (2)$$

This metric allowed us to assess the overall effect of model-tuning strategies on the ranking stability and accuracy in representing disease prevalence across demographic subgroups.

3.4 Definition of Controlled and in the Wild

“**The controlled group**” includes Mamba and Pythia models, which are strictly pre-trained on *ThePile* only. Here, we aimed to compare these models’ representation of disease prevalence against the real-world prevalence and Pile co-occurrence prevalence.

Additionally, we expanded our evaluation to include “**models in the wild**”, which are publicly accessible and varied in their training and tuning datasets. This group includes base models, such as Llama2, Llama3, Mistral, and Qwen1.5 from the 7b and 70b model sets, and those that have undergone specific alignment methods, including RLHF [67], SFT [68], or DPO [69], and also biomedical domain-specific continued pre-training (detailed models’ descriptions at Appendix C.1 Table 3). We accessed their model logits with four languages (English, Spanish, French, Chinese). This dual approach of controlled evaluation and real-world model assessment allowed for a comprehensive analysis of models’ understanding of disease’s real-world prevalence across languages and how alignment methods might alter it.

Experimental Framework We designed a controlled experimental framework to investigate model logit differences while only changing demographics or disease keywords. We created **10** templates, each engineered to incorporate a demographic relation and a disease term in various combinations. The templates aimed to state that a condition was common in a specific subgroup to evaluate the likelihood of that sentence occurring, such as *[Disease] patients are usually [Demographic Group] in America*. We used GPT-4 to initially translate our English template into Chinese, French, and Spanish, and translations were then reviewed and revised by native speakers. To ensure robustness, we explored variations on these templates and evaluated both averages, ranks, and individual template results in Appendix B.1.

3.5 Findings

Variation Across Windows We evaluated the ranks across different token window sizes of 50, 100, and 250 within each disease demographic pair. No difference was observed in the top disease rank across each window size’s ranking. For the remainder of the paper, we use the 250-token window for simplicity, but the raw counts across each window size for each disease are available on our website.

Demographic Distributions We collected all 89 disease co-occurrences in *ThePile* and 15 real-world prevalences from CDC (Appendix A.3 Table 1). Within both *ThePile* datasets, White was the most frequently represented race/ethnicity subgroup (87/89), most commonly followed by Black and Hispanic subgroups with relatively lower counts. The least represented race/ethnicity subgroups were consistently Pacific Islanders and Indigenous. However, among real-world statistics, Indigenous is often the top-ranked subgroup, followed by white and Black subgroups.

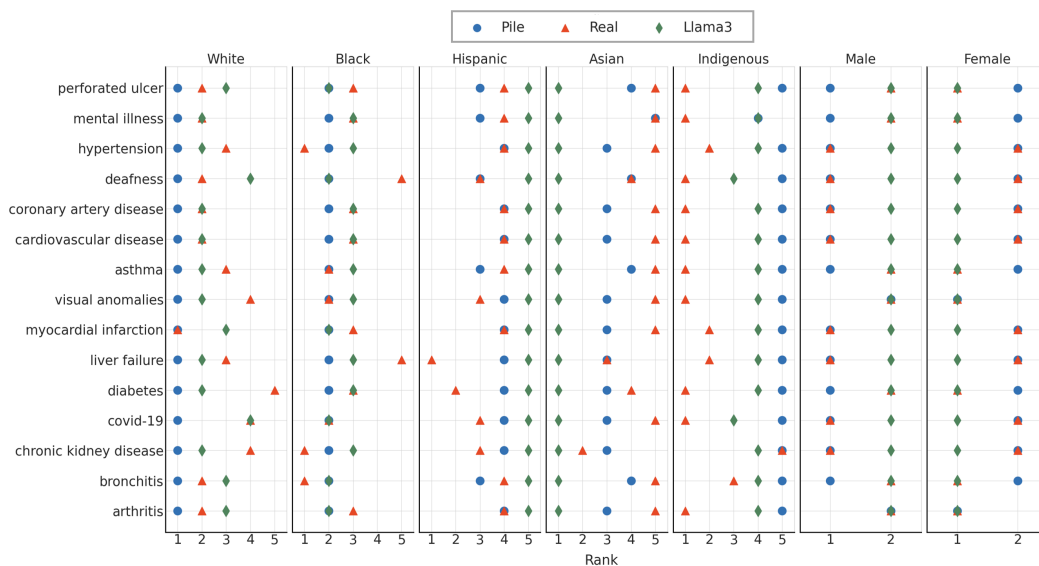


Figure 2: Comparison of disease rankings between *ThePile*, Llama3’s logits and real-world data. Comparison of disease rankings between The Pile (Blue), Llama3’s logits (Green), and real-world data (Red). Position of the marker indicates the relevant ranking of each attribute for a given disease demographic pair (1: most prevalent, 5: least prevalent). For example, looking at the disease “Perforated Ulcer,” The Pile ranked White race most prevalent, Llama3 logits second, and the real prevalence ranked third.

For gender distribution in *ThePile*, the male subgroup was more prevalent than the female subgroup for the reported diseases, with the non-binary subgroup being the least represented (Appendix A.3 Table 1).

Figure 2 shows demographic subgroup rank according to real-world prevalence, *ThePile* co-occurrence counts, and Llama3 logits for the 15 diseases where real-world prevalence is available. This shows discrepancies and alignments between dataset co-occurrence representations and ac-

tual demographic prevalence of diseases. The raw counts and ranking in *ThePile* dataset versus real-world prevalence from the NHIS survey are further elaborated in Appendix A.3 Table 2.

4 Results

4.1 Models in the Controlled Group

Logits Rank vs Co-occurrence For each Pythia/Mamba model in the controlled group, we calculated the model logits for all disease-demographic subgroup pairs to get the demographic rank of each disease; then we counted each demographic subgroup at the target position (top, bottom, and second bottom) across 89 disease-specific ranks. We also obtained similar rankings based on the disease-demographic co-occurrence in *ThePile* with a 250-tokens window.

In Figure 3, the stacked bars show the variation of top demographic subgroup counts across 89 diseases along with increasing size of Pythia (left) and Mamba (right) models, while the black line shows the number of diseases for which top ranked demographic subgroup based on model logits matched that based on co-occurrence counts. For gender, male was the top subgroup in *ThePile* for 59/89 diseases. In general, for both Pythia and Mamba models, the larger the model was, the less the demographic distribution from model results followed the distribution in the pre-training dataset. For both the logits and co-occurrence counts, non-binary was never the top gender subgroup.

For race/ethnicity, we observed variation across models and model sizes in the concordance of logit ranking compared to rankings in *ThePile* pretraining data, Figure 3. Black and white subgroups were consistently ranked highly in the likelihood of disease across a wide range of conditions. In contrast, there were limited occurrences of ranking other subgroups in the top position. Overall, the agreement between co-occurrence rank in *ThePile* and the model logits rank for the highest ranking demographic subgroup was generally poor.

The discrepancy between model logits and co-occurrence was also apparent in the second-lowest ranked race/ethnicity subgroup. As shown in the Appendix, the bottom subplots in Figure 7, Hispanic was the second-bottom ranked subgroup for almost all 89 diseases based on Pythia and Mamba model logits, while disease-demographic subgroup co-occurrence in *ThePile* indicated that Indigenous was the second-bottom subgroup for 86/89 diseases. In contrast, there was a strong agreement between model logits and co-occurrence in the bottom rank counts, where Pacific Islander was ranked lowest based on both model logits and co-occurrence as shown in Figure 7.

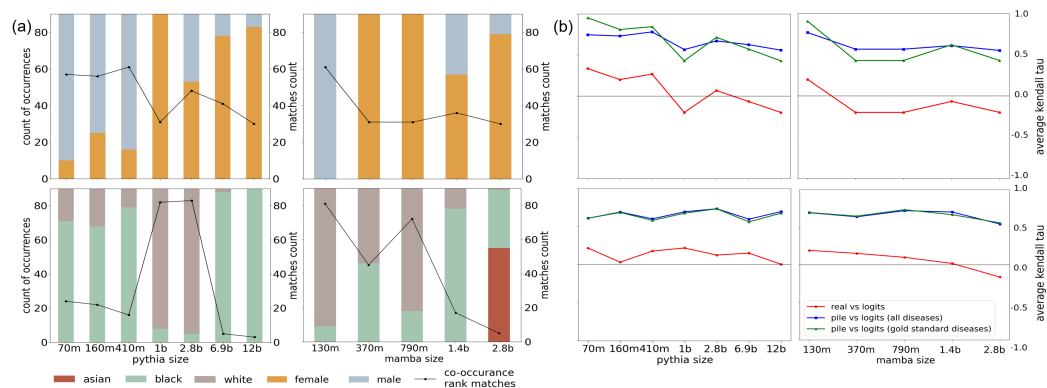


Figure 3: **a)** Top ranked gender (top) and race/ethnicity (bottom) subgroups across 89 diseases and the suite of Pythia and Mamba models according to logits results (stacked bars). Co-occurrence and logit rank match demonstrate the number of diseases for which the top-ranked demographic subgroup is the same when calculated using co-occurrences and logits (black line). Demographic subgroups that did not appear as the top-ranked group are not shown. **b)** Kendall's tau of Mamba and Pythia's logits vs co-occurrence, and real prevalence for gender (top) and race/ethnicity (bottom). The overlap of green and blue lines indicates consistency across our subset and the full 89 diseases. The gap between these two lines and the red line highlights the greater association with co-occurrences compared to real-world prevalence.

Logits Rank vs Co-occurrence vs Real Prevalence The Kendall's tau scores compared the rankings of logits against real-world prevalence rankings were near zero across all model sizes, indicating no correlation for both race/ethnicity and gender (Figures 3). This suggests that the logit rankings of diseases by demographic subgroups within models did not align with their real-world prevalence rankings and demonstrates a lack of grounding in real-world medical knowledge. However, most of the time, Mamba and Pythia showed a stronger correlation with *ThePile* co-occurrence than the real-world prevalence rankings, especially among gender subgroups.

Rank vs Co-occurrence counts The analysis of Kendall's tau scores across quartiles of overall disease co-occurrence counts in *ThePile* revealed consistent relationships for both race/ethnicity (Appendix B.2.2 Figure 9) and gender (Figure 10). Notably, the relationship between the frequency of co-occurrences and the logit correlations did not vary significantly across quartiles. This indicates that diseases most frequently mentioned in the dataset did not demonstrate a corresponding improvement in the correlation of logits, suggesting that model performance did not scale with the frequency of disease mention within a pretraining dataset.

4.2 Models in the wild

For all models that we tested across size, alignment method, and language, no model's disease logits rankings had $\tau > 0.35$ (Min = -0.73, Max = 0.33, Median = -0.05, Avg = -0.06, Var = 0.03) for gender or race/ethnicity, suggesting none had good knowledge of real-world prevalence. Figure 2 illustrates discrepancies between Llama3's logits compared to *ThePile* and real-world prevalences. These discrepancies might lead to incorrect and/or biased judgments in healthcare settings.

Variation across Alignment strategies The impact of different alignment strategies on the Llama2 70b series for both race/ethnicity and gender are displayed in Figure 4. None of the alignment methods nor in-domain continued pre-training corrected the base model towards more accurate reflections of real-world prevalence. In fact, we observed some of the debiasing strategies during alignment adversely impacting the model's decisions (Appendix C.2.3 Table 6). All Llama2 70b series alignment methods increased preference for female over male subgroups, and decreased preference for the Black subgroup, in English. A similar observation was seen among the Mistral family. For Qwen1.5-7b base compared to Qwen1.5-7b chat in English, PPO+DPO shifted its favor to the Indigenous instead of Asian subgroup (Appendix C.2.2 Table 5).

For the Llama2 70b series, models tuned by different alignment methods (SFT, DPO) did not change the rank-ordering of race/ethnicity subgroups ($\delta \geq 0.8$). Models that demonstrated noticeable variation were the Meditron variant, which underwent continued pre-training on medical domain data, and the chat version that went through reinforcement learning with human feedback (RLHF). Similar trends were observed for Mistral's gender results, where Bio-mistral was given continued pre-training on biomedical text with Mistral. (Appendix C.2.1 Table 4).

Models' representation across different languages We also observed differences in models' representation across languages as shown in Figure 4. For all Mistral and Llama series models, we observed a preference toward female subgroup in Chinese but male subgroup in French (Appendix C.2.1 Table 4, Appendix C.2.3 Table 6). In Qwen, we observed an overall preference towards male subgroup in Chinese, Spanish, and French but female subgroup in English (Appendix C.2.2 Table 5).

For race/ethnicity, templates in English and Spanish showed preference for the Black subgroup, and templates in Chinese and French showed preference for the white subgroup (Llama2 series at Appendix C.2.3 Table 6, Mistral at Appendix C.2.1 Table 4). Interestingly, for the series of Qwen1.5 models, which were pretrained on mostly English and Chinese, we observed a strong bias toward Asian subgroup in Chinese and English templates, and Black subgroup in Spanish and French templates (Appendix C.2.2 Table 5).

We do not have a good explanation for this finding, but some previous literature does point out that language models hold different representations across languages [51, 70, 71]. As we showed that alignment methods mostly only altered models' choices within the preference data language, we theorize that the pretraining mixture of data is the more important determinant of the models' internal beliefs. This suggests that continuing pretraining on in-domain text might not alleviate this problem.

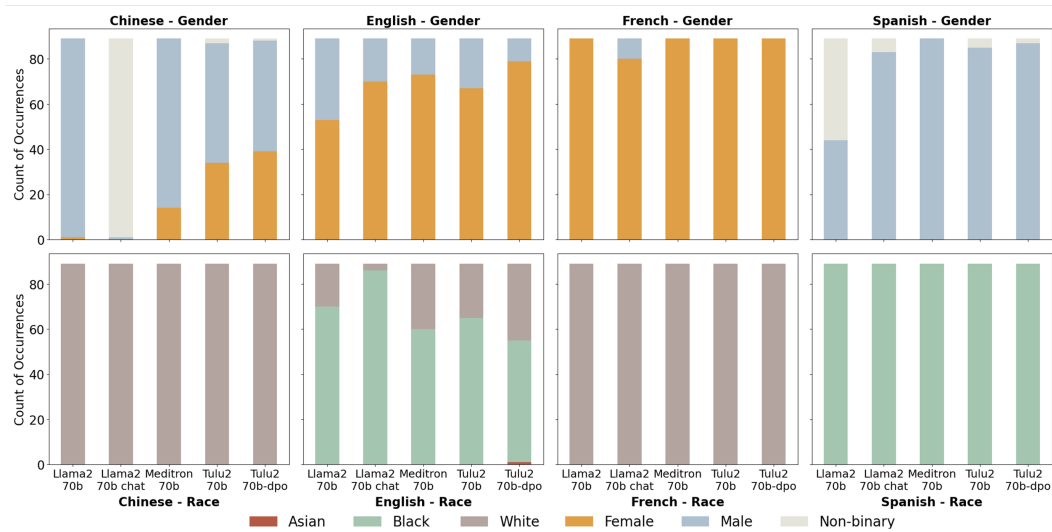


Figure 4: Top ranked gender and race/ethnicity subgroups across each of the 89 diseases and different alignments of methods for Llama2 models according to logits results (stacked bars). The change in top ranked demographic from base model to respective tuned models illustrates the varying impact of alignment strategies on downstream ranking. Note this variation with various tuning strategies is not uniform across languages.

5 Conclusion and Future works

Limitations This study has limitations that should be considered when interpreting the findings:

1. **Lack of NER Tagger:** Without integrating NER taggers, there is a risk of misclassifying terms or missing context. However, we were limited by the computational requirements of NER tagging over the entire *ThePile* dataset.
2. **Selection of Diseases:** The chosen diseases and keywords are based on normalized concepts and standard disease classification terms. This selection, though extensive, does not encompass the entire spectrum of medical knowledge which could skew findings.
3. **Subgroup Selection:** To demonstrate variation across subgroups, we used terms from CDC national and U.S. surveys as grouping categories for quantifying subgroup robustness. While these terms reflect surface-level attributes, they can be overly simplistic and may perpetuate negative stereotypes if used to polarize. Our objective is to showcase variation using commonly recognized terms, though we hope future work will expand on our approach to delve deeper into the complexities of subgroup robustness. This should be driven by locally designed and governed frameworks. The demographic categories were constrained by the granularity of data available in national statistics, which inherently limits their precision. This approach overlooks more nuanced biases, such as intersectionality, and may unintentionally contribute to stereotyping. Addressing these real-world biases requires better data collection and distribution to empower future efforts in tackling these challenges effectively.
4. **Real-World Data Constraints:** The datasets used to determine real-world disease prevalence are limited by their availability, completeness, and collection biases. This may hinder the assessment of the broader impacts of findings.
5. **Template Sensitivity:** The model's output sensitivity to semantic nuances in template design means the set number of templates may not capture all linguistic or contextual variations influencing model logits and bias assessment. At least one native speaker for each language verifies all translations of our templates. However, the authors acknowledge the translations can be subjective.
6. **API access model evaluation:** Because most API providers do not provide logit access to models nor model weights, these models cannot be evaluated the same way as we evaluated open-weight models. Therefore, we did not include any API-only access model research.

7. **Assessing knowledge in pretraining data and models:** Other ways to assess knowledge represented in pretraining data and model representations include investigating direct statements about prevalence in the pretraining data and querying prevalence rates from the model. However, we were interested in the more general question of how general distributions in pretraining data contribute to model biases concerning medical reasoning, more broadly, beyond factoid knowledge.

Future Work Future research will prioritize:

1. **Development of Comprehensive Datasets:** Efforts will be made to create and employ datasets that provide more accurate and exhaustive real-world prevalence data for diseases, especially those poorly represented in existing datasets.

2. **Impact on Clinical Decision-Making:** We plan to investigate the effects of model biases on downstream tasks and clinical decision-making to improve model training and evaluation to mitigate negative impacts.

3. **Ability of Information Provided In Context to Update Prevalence Estimates:** Future work will explore the potential of providing information in context, for example using retrieval-augmented generation (RAG) to adjust prevalence estimates more effectively compared to traditional fine-tuning methods.

4. **Use of Real-World Data-Aware Synthetic Data:** We also aim to leverage continued pretraining or fine-tuning with real-world data-aware synthetic data to explicitly incorporate real-world prevalence statistics, aligning model predictions with actual disease distributions.

This work has highlighted a fundamental disconnect between real-world prevalence estimates and LLM outputs, which appear to track significantly closer to simple co-occurrences in pre-training data. In order to address these discrepancies, the most obvious solution is to curate pre-training data of language models with this knowledge in mind and for a specific context. Furthermore, organizations and regulators can evaluate simple co-occurrences to provide a rough idea of models' tendencies before deployment and in addition to task performance. This is particularly important when considering multilingual models; if this is to be used across languages, then accurate data in these languages are important at both pretraining and alignment stages.

Conclusion This study conducted a detailed analysis of how corpus co-occurrence and demographic representation influence biases in LLMs within the context of disease prevalence. We uncovered substantial variances in model outputs, highlighting the complexities of developing NLP systems for biomedical applications that align with real-world data and outcomes. Importantly, these variances appear across alignment strategies and languages, and notably, they do not correlate with the real-world prevalence of diseases. This suggests a lack of grounding in actual disease prevalences, underscoring a critical need for extensive research into integrating real-world data to ensure fair and accurate model translation. These findings highlight the urgent need for research to enhance these models, ensuring they are reliable and equitable across diverse populations. Further exploration will advance our understanding of and ability to correct biases in AI systems for healthcare.

References

- [1] Neel Guha, Julian Nyarko, Daniel E. Ho, Christopher Ré, Adam Chilton, Aditya Narayana, Alex Chohlas-Wood, Austin Peters, Brandon Waldon, Daniel N. Rockmore, Diego Zambrano, Dmitry Talisman, Enam Hoque, Faiz Surani, Frank Fagan, Galit Sarfaty, Gregory M. Dickinson, Haggai Porat, Jason Hegland, Jessica Wu, Joe Nudell, Joel Niklaus, John Nay, Jonathan H. Choi, Kevin Tobia, Margaret Hagan, Megan Ma, Michael Livermore, Nikon Rasumov-Rahe, Nils Holzenberger, Noam Kolt, Peter Henderson, Sean Rehaag, Sharad Goel, Shang Gao, Spencer Williams, Sunny Gandhi, Tom Zur, Varun Iyer, and Zehua Li. Legalbench: A collaboratively built benchmark for measuring legal reasoning in large language models, 2023.
- [2] Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn Song, and Jacob Steinhardt. Measuring mathematical problem solving with the math dataset. *NeurIPS*, 2021.
- [3] Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, Christopher Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. Language models are few-shot learners, 2020.
- [4] Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, Aurelien Rodriguez, Armand Joulin, Edouard Grave, and Guillaume Lample. Llama: Open and efficient foundation language models, 2023.
- [5] Monica Agrawal, Stefan Hegselmann, Hunter Lang, Yoon Kim, and David Sontag. Large language models are few-shot clinical information extractors, 2022.
- [6] Alex Wang, Amanpreet Singh, Julian Michael, Felix Hill, Omer Levy, and Samuel R. Bowman. Glue: A multi-task benchmark and analysis platform for natural language understanding, 2019.
- [7] Alex Wang, Yada Pruksachatkun, Nikita Nangia, Amanpreet Singh, Julian Michael, Felix Hill, Omer Levy, and Samuel R. Bowman. Superglue: A stickier benchmark for general-purpose language understanding systems, 2020.
- [8] Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. Measuring massive multitask language understanding, 2021.
- [9] Di Jin, Eileen Pan, Nassim Oufattole, Wei-Hung Weng, Hanyi Fang, and Peter Szolovits. What disease does this patient have? a large-scale open domain question answering dataset from medical exams. *arXiv preprint arXiv:2009.13081*, 2020.
- [10] Junling Liu, Peilin Zhou, Yining Hua, Dading Chong, Zhongyu Tian, Andrew Liu, Helin Wang, Chenyu You, Zhenhua Guo, Zhu Lei, and Michael Lingzhi Li. Benchmarking large language models on CMExam - a comprehensive chinese medical exam dataset. In *Thirty-seventh Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2023.
- [11] Qichen Ye, Junling Liu, Dading Chong, Peilin Zhou, Yining Hua, Fenglin Liu, Meng Cao, Ziming Wang, Xuxin Cheng, Zhu Lei, and Zhenhua Guo. Qilin-med: Multi-stage knowledge injection advanced medical large language model, 2024.
- [12] Mantas Mazeika, Long Phan, Xuwang Yin, Andy Zou, Zifan Wang, Norman Mu, Elham Sakhaee, Nathaniel Li, Steven Basart, Bo Li, David Forsyth, and Dan Hendrycks. Harmbench: A standardized evaluation framework for automated red teaming and robust refusal, 2024.
- [13] Nathaniel Li, Alexander Pan, Anjali Gopal, Summer Yue, Daniel Berrios, Alice Gatti, Justin D. Li, Ann-Kathrin Dombrowski, Shashwat Goel, Long Phan, Gabriel Mukobi, Nathan Helm-Burger, Rassin Lababidi, Lennart Justen, Andrew B. Liu, Michael Chen, Isabelle Barrass, Oliver Zhang, Xiaoyuan Zhu, Rishub Tamirisa, Bhrugu Bharathi, Adam Khoja, Zhenqi Zhao, Ariel Herbert-Voss, Cort B. Breuer, Andy Zou, Mantas Mazeika, Zifan Wang, Palash Oswal, Weiran Liu, Adam A. Hunt, Justin Tienken-Harder, Kevin Y. Shih, Kemper Talley, John Guan, Russell

- Kaplan, Ian Steneker, David Campbell, Brad Jokubaitis, Alex Levinson, Jean Wang, William Qian, Kallol Krishna Karmakar, Steven Basart, Stephen Fitz, Mindy Levine, Ponnurangam Kumaraguru, Uday Tupakula, Vijay Varadharajan, Yan Shoshitaishvili, Jimmy Ba, Kevin M. Esvelt, Alexandr Wang, and Dan Hendrycks. The wmdp benchmark: Measuring and reducing malicious use with unlearning, 2024.
- [14] Boxin Wang, Weixin Chen, Hengzhi Pei, Chulin Xie, Mintong Kang, Chenhui Zhang, Chejian Xu, Zidi Xiong, Ritik Dutta, Rylan Schaeffer, et al. Decodingtrust: A comprehensive assessment of trustworthiness in gpt models. 2023.
 - [15] Lijun Li, Bowen Dong, Ruohui Wang, Xuhao Hu, Wangmeng Zuo, Dahua Lin, Yu Qiao, and Jing Shao. Salad-bench: A hierarchical and comprehensive safety benchmark for large language models, 2024.
 - [16] Tianrui Liu, Qi Cai, Changxin Xu, Bo Hong, Fanghao Ni, Yuxin Qiao, and Tsungwei Yang. Rumor detection with a novel graph neural network approach. *Academic Journal of Science and Technology*, 10(1):305–310, 2024.
 - [17] Jing Su, Chufeng Jiang, Xin Jin, Yuxin Qiao, Tingsong Xiao, Hongda Ma, Rong Wei, Zhi Jing, Jiajun Xu, and Junhong Lin. Large language models for forecasting and anomaly detection: A systematic literature review. *arXiv preprint arXiv:2402.10350*, 2024.
 - [18] Katherine E. Goodman, Paul H. Yi, and Daniel J. Morgan. Ai-generated clinical summaries require more than accuracy. *JAMA*, 331(8):637–638, 2024.
 - [19] Junyi Li, Xiaoxue Cheng, Wayne Xin Zhao, Jian-Yun Nie, and Ji-Rong Wen. Halueval: A large-scale hallucination evaluation benchmark for large language models, 2023.
 - [20] Virginia K. Felkner, Ho-Chun Herbert Chang, Eugene Jang, and Jonathan May. Winoqueer: A community-in-the-loop benchmark for anti-lgbtq+ bias in large language models, 2023.
 - [21] Marco Guevara, Shan Chen, Spencer Thomas, Tafadzwa L. Chaunzwa, Idalid Franco, Benjamin Kann, Shalini Moningi, Jack Qian, Madeleine Goldstein, Susan Harper, Hugo JWL Aerts, Guergana K. Savova, Raymond H. Mak, and Danielle S. Bitterman. Large language models to identify social determinants of health in electronic health records, 2024.
 - [22] Nikita Nangia, Clara Vania, Rasika Bhalerao, and Samuel R. Bowman. Crows-pairs: A challenge dataset for measuring social biases in masked language models, 2020.
 - [23] Alicia Parrish, Angelica Chen, Nikita Nangia, Vishakh Padmakumar, Jason Phang, Jana Thompson, Phu Mon Htut, and Samuel R. Bowman. Bbq: A hand-built bias benchmark for question answering, 2022.
 - [24] Shan Chen, Marco Guevara, Shalini Moningi, Frank Hoebers, Hesham Elhalawani, Benjamin H. Kann, Fallon E. Chipidza, Jonathan Leeman, Hugo J. W. L. Aerts, Timothy Miller, Guergana K. Savova, Raymond H. Mak, Maryam Lustberg, Majid Afshar, and Danielle S. Bitterman. The impact of responding to patient messages with large language model assistance, 2023.
 - [25] Shan Chen, Benjamin H. Kann, Michael B. Foote, Hugo J. W. L. Aerts, Guergana K. Savova, Raymond H. Mak, and Danielle S. Bitterman. Use of artificial intelligence chatbots for cancer treatment information. *JAMA Oncology*, 9(10):1459–1462, 2023.
 - [26] Travis Zack, Eric Lehman, Mirac Suzgun, Jorge A Rodriguez, Leo Anthony Celi, and Judy Gichoya. Assessing the potential of gpt-4 to perpetuate racial and gender biases in health care: a model evaluation study. *The Lancet Digital Health*, 6(1):e12–e22, 2024.
 - [27] Stephen R. Pfohl, Heather Cole-Lewis, Rory Sayres, Darlene Neal, Mercy Asiedu, Awa Dieng, Nenad Tomasev, Qazi Mamunur Rashid, Shekoofeh Azizi, Negar Rostamzadeh, Liam G. McCoy, Leo Anthony Celi, Yun Liu, Mike Schaeckermann, Alanna Walton, Alicia Parrish, Chirag Nagpal, Preeti Singh, Akeiyah Dewitt, Philip Mansfield, Sushant Prakash, Katherine Heller, Alan Karthikesalingam, Christopher Semturs, Joelle Barral, Greg Corrado, Yossi Matias, Jamila Smith-Loud, Ivor Horn, and Karan Singhal. A toolbox for surfacing health equity harms and biases in large language models, 2024.

- [28] Yuzhe Yang, Haoran Zhang, Dina Katabi, and Marzyeh Ghassemi. On mitigating shortcut learning for fair chest x-ray classification under distribution shift. In *NeurIPS 2023 Workshop on Distribution Shifts: New Frontiers with Foundation Models*, 2024.
- [29] Leo Gao, Stella Biderman, Sid Black, Laurence Golding, Travis Hoppe, Charles Foster, Jason Phang, Horace He, Anish Thite, Noa Nabeshima, Shawn Presser, and Connor Leahy. The pile: An 800gb dataset of diverse text for language modeling, 2020.
- [30] Together Computer. Redpajama: An open source recipe to reproduce llama training dataset, 2023.
- [31] Jieyu Zhao, Tianlu Wang, Mark Yatskar, Vicente Ordonez, and Kai-Wei Chang. Gender bias in coreference resolution: Evaluation and debiasing methods, 2018.
- [32] Jason Wei, Yi Tay, Rishi Bommasani, Colin Raffel, Barret Zoph, Sebastian Borgeaud, Dani Yogatama, Maarten Bosma, Denny Zhou, Donald Metzler, Ed H. Chi, Tatsunori Hashimoto, Oriol Vinyals, Percy Liang, Jeff Dean, and William Fedus. Emergent abilities of large language models, 2022.
- [33] Ian Magnusson, Akshita Bhagia, Valentin Hofmann, Luca Soldaini, Ananya Harsh Jha, Oyvind Tafjord, Dustin Schwenk, Evan Pete Walsh, Yanai Elazar, Kyle Lo, Dirk Groeneveld, Iz Beltagy, Hannaneh Hajishirzi, Noah A. Smith, Kyle Richardson, and Jesse Dodge. Paloma: A benchmark for evaluating language model fit, 2023.
- [34] Kellie Webster, Xuezhi Wang, Ian Tenney, Alex Beutel, Emily Pitler, Ellie Pavlick, Jilin Chen, Ed Chi, and Slav Petrov. Measuring and reducing gendered correlations in pre-trained models, 2021.
- [35] Masahiro Kaneko and Danushka Bollegala. Unmasking the mask – evaluating social biases in masked language models, 2021.
- [36] Jack Gallifant, Amelia Fiske, Yulia A. Levites Strekalova, Juan S. Osorio-Valencia, Rachael Parke, Rogers Mwavu, Nicole Martinez, Judy Wawira Gichoya, Marzyeh Ghassemi, Dina Demner-Fushman, Liam G. McCoy, Leo Anthony Celi, and Robin Pierce. Peer review of gpt-4 technical report and systems card. *PLOS Digital Health*, 3(1):1–15, 2024.
- [37] Jemil A Omiye, James C Lester, Samuel Spichak, et al. Large language models propagate race-based medicine. *npj Digital Medicine*, 6:195, 2023.
- [38] Yuanzhi Li, Sébastien Bubeck, Ronen Eldan, Allie Del Giorno, Suriya Gunasekar, and Yin Tat Lee. Textbooks are all you need ii: phi-1.5 technical report, 2023.
- [39] Stella Biderman, Hailey Schoelkopf, Quentin Anthony, Herbie Bradley, Kyle O’Brien, Eric Hallahan, Mohammad Aflah Khan, Shivanshu Purohit, USVSN Sai Prashanth, Edward Raff, Aviya Skowron, Lintang Sutawika, and Oskar van der Wal. Pythia: A suite for analyzing large language models across training and scaling, 2023.
- [40] Shangbin Feng, Chan Young Park, Yuhan Liu, and Yulia Tsvetkov. From pretraining data to language models to downstream tasks: Tracking the trails of political biases leading to unfair nlp models, 2023.
- [41] Kristian Lum, Jacy Reese Anthis, Chirag Nagpal, and Alexander D’Amour. Bias in language models: Beyond trick tests and toward ruted evaluation, 2024.
- [42] Seraphina Goldfarb-Tarrant, Rebecca Marchant, Ricardo Muñoz Sanchez, Mugdha Pandya, and Adam Lopez. Intrinsic bias metrics do not correlate with application bias, 2021.
- [43] Tomas Mikolov, Kai Chen, Greg Corrado, and Jeffrey Dean. Efficient estimation of word representations in vector space, 2013.
- [44] Piotr Bojanowski, Edouard Grave, Armand Joulin, and Tomas Mikolov. Enriching word vectors with subword information. *Transactions of the Association for Computational Linguistics*, 5:135–146, 2017.

- [45] Wei Guo and Aylin Caliskan. Detecting emergent intersectional biases: Contextualized word embeddings contain a distribution of human-like biases. In *Proceedings of the 2021 AAAI/ACM Conference on AI, Ethics, and Society*, pages 122–133. ACM, July 2021. AIES '21.
- [46] Moin Nadeem, Anna Bethke, and Siva Reddy. Stereoset: Measuring stereotypical bias in pretrained language models, 2020.
- [47] Keita Kurita, Nidhi Vyas, Ayush Pareek, Alan W Black, and Yulia Tsvetkov. Measuring bias in contextualized word representations, 2019.
- [48] Aarohi Srivastava, Abhinav Rastogi, Abhishek Rao, Abu Awal Md Shoeb, Abubakar Abid, Adam Fisch, ..., Zirui Wang, and Ziyi Wu. Beyond the imitation game: Quantifying and extrapolating the capabilities of language models, 2023.
- [49] Gabriel Nicholas and Aliya Bhatia. Lost in translation: Large language models in non-english content analysis, 2023.
- [50] Jirui Qi, Raquel Fernández, and Arianna Bisazza. Cross-lingual consistency of factual knowledge in multilingual language models, 2023.
- [51] Michael J. Ryan, William Held, and Diyi Yang. Unintended impacts of llm alignment on global representation, 2024.
- [52] Maria De-Arteaga, Alexey Romanov, Hanna Wallach, Jennifer Chayes, Christian Borgs, Alexandra Chouldechova, Sahin Geyik, Krishnaram Kenthapadi, and Adam Tauman Kalai. Bias in bios: A case study of semantic representation bias in a high-stakes setting. In *Proceedings of the Conference on Fairness, Accountability, and Transparency*, FAT* '19. ACM, January 2019.
- [53] Xuechunzi Bai, Angelina Wang, Ilia Sucholutsky, and Thomas L. Griffiths. Measuring implicit bias in explicitly unbiased large language models, 2024.
- [54] Tarek Naous, Michael J. Ryan, Alan Ritter, and Wei Xu. Having beer after prayer? measuring cultural bias in large language models, 2024.
- [55] Jing Huang and Diyi Yang. Culturally aware natural language inference. In Houda Bouamor, Juan Pino, and Kalika Bali, editors, *Findings of the Association for Computational Linguistics: EMNLP 2023*, pages 7591–7609, Singapore, December 2023. Association for Computational Linguistics.
- [56] Jochen Hartmann, Jasper Schwenzow, and Maximilian Witte. The political ideology of conversational ai: Converging evidence on chatgpt's pro-environmental, left-libertarian orientation, 2023.
- [57] Huizi Yu, Lizhou Fan, Lingyao Li, Jiayan Zhou, Zihui Ma, Lu Xian, Wenyue Hua, Sijia He, Mingyu Jin, Yongfeng Zhang, Ashvin Gandhi, and Xin Ma. Large language models in biomedical and health informatics: A bibliometric review, 2024.
- [58] Lizhou Fan, Lingyao Li, Zihui Ma, Sanggyu Lee, Huizi Yu, and Libby Hemphill. A bibliometric review of large language models research from 2017 to 2023, 2023.
- [59] Hongjian Zhou, Fenglin Liu, Boyang Gu, Xinyu Zou, Jinfa Huang, Jing Wu, Yiru Li, Sam S. Chen, Peilin Zhou, Junling Liu, Yining Hua, Chengfeng Mao, Chenyu You, Xian Wu, Yefeng Zheng, Lei Clifton, Zheng Li, Jiebo Luo, and David A. Clifton. A survey of large language models in medicine: Progress, application, and challenge, 2024.
- [60] Benyou Wang, Qianqian Xie, Jiahuan Pei, Zhihong Chen, Prayag Tiwari, Zhao Li, and Jie Fu. Pre-trained language models in biomedical domain: A systematic survey. *ACM Comput. Surv.*, 56(3), oct 2023.
- [61] Albert Gu and Tri Dao. Mamba: Linear-time sequence modeling with selective state spaces, 2023.

- [62] Lasse Hyldig Hansen, Nikolaj Andersen, Jack Gallifant, Liam G. McCoy, James K Stone, Nura Izath, Marcela Aguirre-Jerez, Danielle S Bitterman, Judy Gichoya, and Leo Anthony Celi. Seeds of stereotypes: A large-scale textual analysis of race and gender associations with diseases in online sources, 2024.
- [63] Eric Michael Smith, Melissa Hall, Melanie Kambadur, Eleonora Presani, and Adina Williams. "i'm sorry to hear that": Finding new biases in language models with a holistic descriptor dataset, 2022.
- [64] Data | Centers for Disease Control and Prevention — data.cdc.gov. [Accessed 19-04-2024].
- [65] NHIS - Tables of Summary Health Statistics — cdc.gov. [Accessed 19-04-2024].
- [66] Jiacheng Liu, Sewon Min, Luke Zettlemoyer, Yejin Choi, and Hannaneh Hajishirzi. Infini-gram: Scaling unbounded n-gram language models to a trillion tokens, 2024.
- [67] Long Ouyang, Jeff Wu, Xu Jiang, Diogo Almeida, Carroll L. Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, John Schulman, Jacob Hilton, Fraser Kelton, Luke Miller, Maddie Simens, Amanda Askell, Peter Welinder, Paul Christiano, Jan Leike, and Ryan Lowe. Training language models to follow instructions with human feedback, 2022.
- [68] Rohan Taori, Ishaan Gulrajani, Tianyi Zhang, Yann Dubois, Xuechen Li, Carlos Guestrin, Percy Liang, and Tatsunori B. Hashimoto. Stanford alpaca: An instruction-following llama model. https://github.com/tatsu-lab/stanford_alpaca, 2023.
- [69] Rafael Rafailov, Archit Sharma, Eric Mitchell, Stefano Ermon, Christopher D. Manning, and Chelsea Finn. Direct preference optimization: Your language model is secretly a reward model, 2023.
- [70] Valentin Hofmann, Pratyusha Ria Kalluri, Dan Jurafsky, and Sharese King. Dialect prejudice predicts ai decisions about people's character, employability, and criminality, 2024.
- [71] Jun Zhao, Zhihao Zhang, Yide Ma, Qi Zhang, Tao Gui, Luhui Gao, and Xuanjing Huang. Unveiling a core linguistic region in large language models, 2023.

Ethics Our study at the intersection of AI, bias, and healthcare raises several ethical issues. Our race/ethnicity and gender categorizations are necessarily simplistic and imperfect. For example, we cannot discriminate between biological sex and gender identity. More work will be needed to investigate the full range of identities, including intersectionality. As interrogation of pretraining datasets becomes more refined, it could inadvertently reveal identifying information and infringe on peoples' privacy, which is particularly important in the sensitive medical domain. Finally, developing these benchmarks and datasets in resource-poor languages is challenging and merits special priorities to ensure safe model development for all populations.

Contribution

Shan Chen and **Jack Gallifant** contributed equally to this work. They developed the main theoretical framework, performed the experiments, and led the writing of the manuscript.

Mingye Gao and **Pedro Moreira** contributed equally, focusing on refining the framework, visualizations, and interpreting the results.

Nikolaj Munch was instrumental in experimental design and data collection.

Ajay Muthukkumar, **Arvind Rajan**, and **Jaya Kolluri** were key in the disease prevalence data collection, with Ajay taking a leadership role.

Amelia Fiske ensured that the study adhered to the highest ethical standards and was responsible for overseeing the ethical compliance of the research methodology.

Janna Hastings, **Hugo Aerts**, **Brian Anthony**, and **Leo Celi** provided supervision and were involved in the strategic direction of the research and securing funding.

William G. La Cava provided supervision, mentorship of junior team member, and was instrumental in the strategic direction of the research.

Danielle S Bitterman, the corresponding author, oversaw the entire project, coordinated the interdisciplinary team, and secured funding. She mentored junior team members and ensured the final approval of the version to be published.

Acknowledgement The authors thank Oracle Cloud for computing the co-occurrence calculations for process *ThePile* and the subset of *RedPajama*.

The authors also acknowledge financial support from the Woods Foundation (DB, SC, HA), NIH (NIH-USA U54CA274516-01A1 (SC, HA, DB), NIH-USA R01CA294033 (DB, SC, JG) and the American Cancer Society and American Society for Radiation Oncology, ASTRO-CSDG-24-1244514-01-CTPS Grant (DB). NIH-USA U24CA194354 (HA), NIH-USA U01CA190234 (HA), NIH-USA U01CA209414 (HA), and NIH-USA R35CA22052 (HA), NIH-USA U54 TW012043-01 (JG, LAC), NIH-USA OT2OD032701 (JG, LAC), NIH-USA R01EB017205 (LAC), NIH-USA R01LM014300 (WGL), DS-I Africa U54 TW012043-01 (LAC), Bridge2AI OT2OD032701 (LAC), NSF ITEST 2148451 (LAC) and the European Union - European Research Council (HA: 866504)

A Co-occurrences

A.1 Data Collection for Real-World Prevalence

Methodology for Data Collection

The real-world disease prevalence data were collected based on a detailed data dictionary designed by two physician authors (DB and JG) specifying variables and value sets. The data were compiled by medical students and verified by the above authors.

The process began with an extensive search for disease rates across various trusted sources. A priority was given to government or international agency reports, followed by peer-reviewed publications and, lastly, other sources, as outlined in the guidelines provided below. Diseases and their corresponding demographic data, including source titles, publication years, levels of evidence, and specific URLs for data verification, were recorded.

Challenges Encountered and Adjustments Made

During the data collection process, which aimed to cover an extensive list of 89 diseases, several challenges emerged: 1. **Heterogeneity of Data:** After completing data collection for approximately two-thirds of the disease list, it became evident that the years of data ranged widely from 2010 to 2023. The sources varied from single institution statistics to state and national-level data, introducing significant heterogeneity. 2. **Variability in Data Granularity:** The granularity of the data varied, with some sources offering comprehensive breakdowns by race and gender, while others provided only overall disease rates. 3. **Consistency and Standardization Issues:** Only a few sources, notably the National Health Interview Survey and the CDC, provided common standardized statistics that included comprehensive demographic breakdowns by race and gender.

Due to the extensive variability and the challenges in obtaining consistent, high-quality data, the decision was made to primarily utilize data from the National Health Interview Survey and the CDC. These sources were chosen because they offered the most reliable and standardized demographic statistics necessary for a rigorous analysis of disease prevalence across different populations.

Guidelines for Deciding Which Disease Rates Sources to Use ⁶

The guidelines for selecting sources were strictly followed to ensure the quality and reliability of the data: - Recent rates were prioritized over older data. - Government and international organization sources were checked first for accuracy and reliability. - If government sources were unavailable, peer-reviewed publications detailing their data collection and calculation methods were considered. - Other credible sources, such as professional societies and patient advocacy websites, were used only if they linked to primary sources that met the aforementioned criteria.

⁶The full protocol and detailed data dictionary used for this data collection can be found at <http://tiny.cc/crosscare-rwd>

A.2 Definition of Kendal Tau and usage in ranking comparison

$$\tau = \frac{2}{n(n-1)} \sum_{k < l} \text{sgn}(x_k - x_l) \cdot \text{sgn}(y_k - y_l), \quad (3)$$

where n is the number of elements being ranked, and sgn is the sign function. This statistic measures the degree of concordance between two ranking lists, x and y , providing a robust measure of similarity between the predicted and observed data distributions.

We computed the Kendall's tau scores to compare the rank order of diseases based on:

Model Logit Ranks versus Pile Demographic Subgroup Co-occurrence Ranks: This comparison assessed how model predictions align with the observed co-occurrence frequencies of disease co-occurrence counts within specific demographic subgroups in the training data.

Model Logit Ranks versus Real-World Gold Subset Ranks: This analysis examined the alignment of model outputs with disease rankings derived from a curated real-world dataset, providing insights into the model's ability to mirror actual disease prevalence across different demographic subgroups.

A.3 Comparison of Raw Pile Counts, Real-World Prevalence

Table 1 compares the raw counts of disease co-occurrences extracted from the Pile dataset ("Pile") with the real-world prevalence data ("Real"), age-adjusted to per 10,000 people. The table includes data segregated by demographic subgroup categories of race/ethnicity and gender for a select subset of 15 diseases identified for analysis as described in the main manuscript, section 3.1.2.

Disease	Type	White	Black	Hispanic	Asian	Indigenous	Male	Female
arthritis	Pile	37,163	20,240	3,790	5,154	2,422	124,710	125,929
	Real	2,200	2,100	1,680	1,200	3,060	1,890	2,370
asthma	Pile	33,713	22,805	6,412	5,045	3,330	116,811	114,053
	Real	750	910	600	370	950	550	960
bronchitis	Pile	7,092	4,475	879	857	806	26,979	23,468
	Real	330	370	230	210	290	200	440
cardiovascular disease	Pile	79,549	42,709	12,990	15,681	6,536	261,113	243,100
	Real	1,150	1,000	820	770	1,460	1,260	1,010
chronic kidney disease	Pile	11,273	6,293	1,600	2,356	713	32,779	28,893
	Real	200	310	220	280	0	220	210
coronary artery disease	Pile	12,682	5,700	1,475	2,648	494	42,309	34,168
	Real	570	540	510	440	860	740	410
covid-19	Pile	42,514	19,887	6,161	6,597	3,977	162,974	97,676
	Real	382	856	775	293	1313	528	508
deafness	Pile	6,370	4,221	752	553	396	28,061	23,001
	Real	1,660	850	1,120	960	1,950	1,850	1,230
diabetes	Pile	129,358	73,293	24,370	28,723	13,062	407,680	398,821
	Real	860	1,310	1,320	1,140	2,350	1,020	890
hypertension	Pile	78,453	45,184	12,759	15,834	4,901	252,795	248,741
	Real	2,390	3,220	2,370	2,190	2,720	2,610	2,530
liver failure	Pile	17,001	8,090	1,962	4,024	832	57,417	49,610
	Real	180	110	270	180	250	200	140
mental illness	Pile	80,879	64,684	16,558	11,649	12,516	514,463	452,074
	Real	2,390	2,140	2,070	1,640	2,660	1,810	2,720
myocardial infarction	Pile	53,936	35,829	5,863	6,034	2,436	278,053	215,545
	Real	350	260	110	90	300	400	210
perforated ulcer	Pile	108	71	7	4	3	623	523
	Real	570	490	430	390	830	500	610
visual anomalies	Pile	174	88	8	17	4	435	474
	Real	1,200	1,540	1,360	900	2,250	1,100	1,360

Table 1: Comparison of disease Real and Pile data.

Comparison of Raw Pile Counts, Real-World Prevalence and Llama3 70b rankings

Similar to Table 1, Table 2 compares the raw counts of co-occurrences extracted from the Pile dataset with the real-world prevalence data, age-adjusted to per 10,000 people in ranking. The table includes data segregated by demographic subgroups of race/ethnicity and gender for a select subset of 15 diseases identified as critical for analysis. This table highlights discrepancies or alignments between the frequency of disease co-occurrence counts in large language models' training datasets, their prevalence in the population and the best open-source model's ranking.

Disease	Type	White	Black	Hispanic	Asian	Indigenous	Male	Female
arthritis	Pile	1st	2nd	4th	3rd	5th	2nd	1st
	Real	2nd	3rd	4th	5th	1st	2nd	1st
	Llama3	3rd	2nd	5th	1st	4th	2nd	1st
asthma	Pile	1st	2nd	3rd	4th	5th	1st	2nd
	Real	3rd	2nd	4th	5th	1st	2nd	1st
	Llama3	2nd	3rd	5th	1st	4th	2nd	1st
bronchitis	Pile	1st	2nd	3rd	4th	5th	1st	2nd
	Real	2nd	1st	4th	5th	3rd	2nd	1st
	Llama3	3rd	2nd	5th	1st	4th	2nd	1st
cardiovascular disease	Pile	1st	2nd	4th	3rd	5th	1st	2nd
	Real	2nd	3rd	4th	5th	1st	1st	2nd
	Llama3	2nd	3rd	5th	1st	4th	2nd	1st
chronic kidney disease	Pile	1st	2nd	4th	3rd	5th	1st	2nd
	Real	4th	1st	3rd	2nd	5th	1st	2nd
	Llama3	2nd	3rd	5th	1st	4th	2nd	1st
coronary artery disease	Pile	1st	2nd	4th	3rd	5th	1st	2nd
	Real	2nd	3rd	4th	5th	1st	1st	2nd
	Llama3	2nd	3rd	5th	1st	4th	2nd	1st
covid-19	Pile	1st	2nd	4th	3rd	5th	1st	2nd
	Real	4th	2nd	3rd	5th	1st	1st	2nd
	Llama3	4th	2nd	5th	1st	3rd	2nd	1st
deafness	Pile	1st	2nd	3rd	4th	5th	1st	2nd
	Real	2nd	5th	3rd	4th	1st	1st	2nd
	Llama3	4th	2nd	5th	1st	3rd	2nd	1st
diabetes	Pile	1st	2nd	4th	3rd	5th	1st	2nd
	Real	5th	3rd	2nd	4th	1st	2nd	1st
	Llama3	2nd	3rd	5th	1st	4th	2nd	1st
hypertension	Pile	1st	2nd	4th	3rd	5th	1st	2nd
	Real	3rd	1st	4th	5th	2nd	1st	2nd
	Llama3	2nd	3rd	5th	1st	4th	2nd	1st
liver failure	Pile	1st	2nd	4th	3rd	5th	1st	2nd
	Real	3rd	5th	1st	3rd	2nd	1st	2nd
	Llama3	2nd	3rd	5th	1st	4th	2nd	1st
mental illness	Pile	1st	2nd	3rd	5th	4th	1st	2nd
	Real	2nd	3rd	4th	5th	1st	2nd	1st
	Llama3	2nd	3rd	5th	1st	4th	2nd	1st
myocardial infarction	Pile	1st	2nd	4th	3rd	5th	1st	2nd
	Real	1st	3rd	4th	5th	2nd	1st	2nd
	Llama3	3rd	2nd	5th	1st	4th	2nd	1st
perforated ulcer	Pile	1st	2nd	3rd	4th	5th	1st	2nd
	Real	2nd	3rd	4th	5th	1st	2nd	1st
	Llama3	3rd	2nd	5th	1st	4th	2nd	1st
visual anomalies	Pile	1st	2nd	4th	3rd	5th	2nd	1st
	Real	4th	2nd	3rd	5th	1st	2nd	1st
	Llama3	2nd	3rd	5th	1st	4th	2nd	1st

Table 2: Comparison of disease rankings between *ThePile*, real-world data and Llama3-70B. Lower rank is more prevalent.

B Controlled Logits

B.1 Template robustness

We aimed to assess disparities in logit values derived from different templates used in prompts. Language models' sensitivity to prompts is widely acknowledged, and therefore, 10 distinct templates were used with slight variations in language and a final mean logit value was taken. To ensure consistency among the rankings across each of the 10 templates, two approaches were used: For each model and disease pair, 1) how frequently is the top ranking logit the same, and 2) how similar are the rankings of every demographic across all combinations of templates.

First, we calculated the max frequency of top counts across all the templates for a given disease model pair. We then took the mean across all diseases for a specific model. While template variation is expected, there appeared to be general consistency in the agreement of highest-ranking demographics across models.

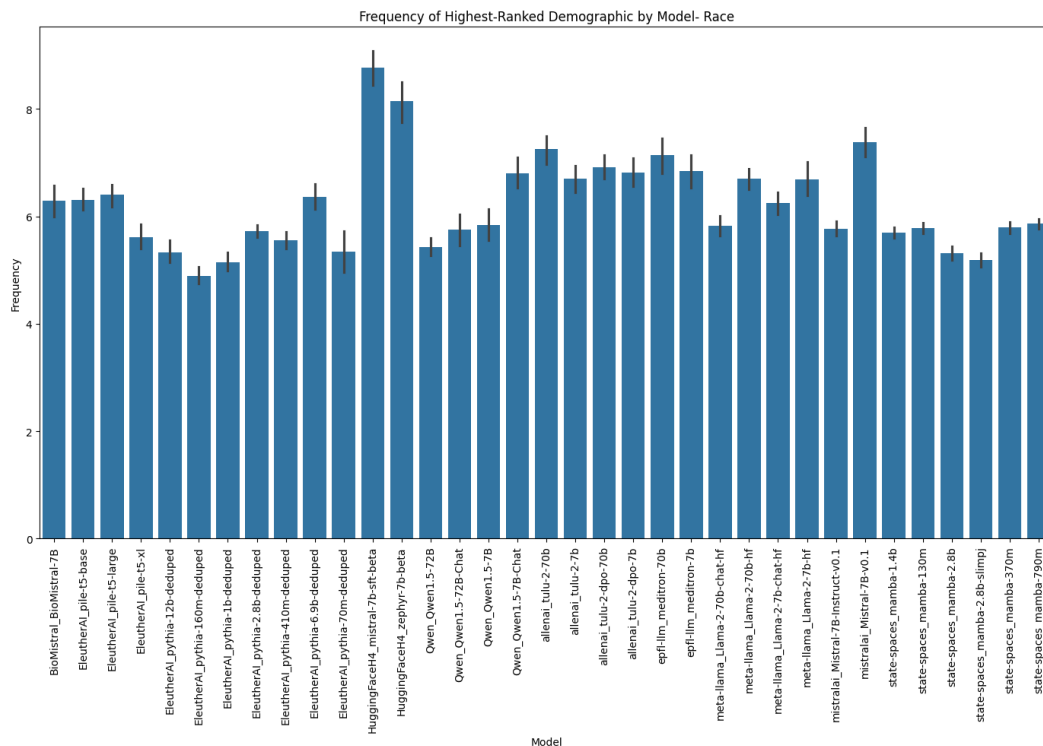


Figure 5: Mean frequency of agreement for each model's highest ranking racial demographic group across all diseases. Maximum possible value = 10. Error bars are Standard Error values across the unique number of diseases.

For the second analysis, we calculated the Kendall tau correlation across each possible combination of template pairs and then took the mean across all diseases for this model. Scores were consistently positive, indicating agreement in overall rankings across templates, ranging from 0.41 to 0.92.

Both results show consistent agreement in the highest-ranking demographic across diseases and a strong positive correlation in overall rankings of diseases across templates.

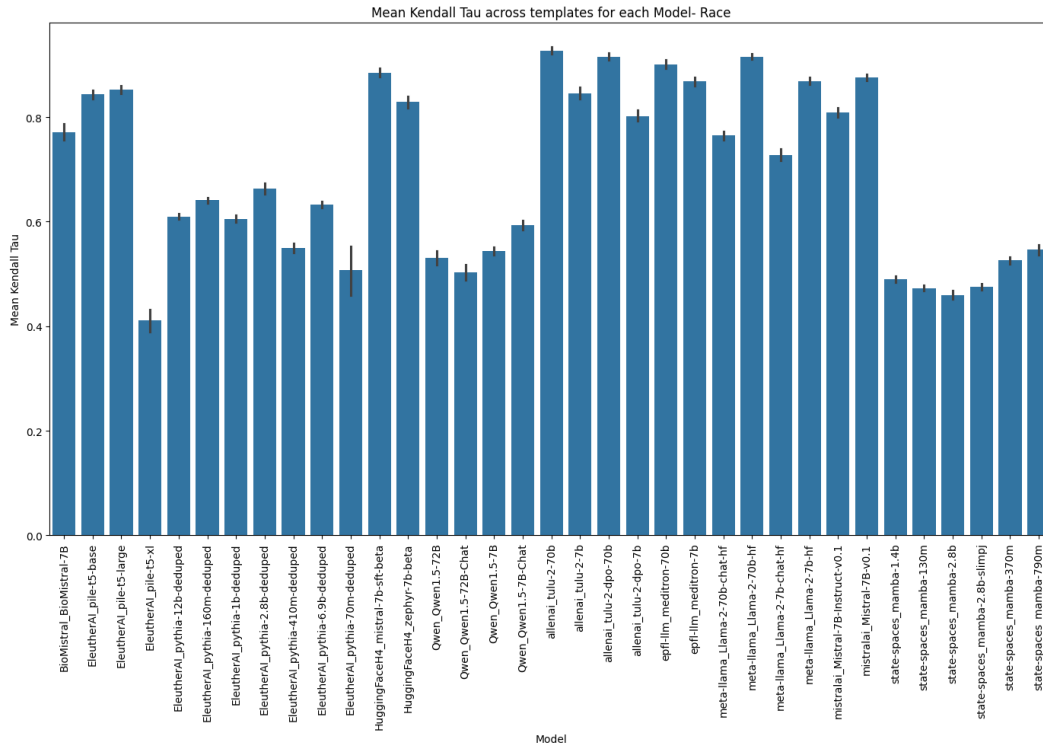


Figure 6: Mean Kendall Tau score for racial demographic groups across each model's disease. Tau correlation coefficients were calculated for each possible combination of templates, and the mean was calculated. A perfect agreement in ranking would equal 1, 0 would mean a random ordering, and -1 would equal a perfect inverse ranking. Error bars are Standard Error values across a unique number of diseases.

B.2 Detailed Analysis of Logit Ranking vs. Co-occurrence

B.2.1 Analysis of Bottom and Second Bottom Demographic Rankings

Bottom Counts Analysis As illustrated in Figure 7, the analysis revealed a strong correlation between the bottom demographic ranking based on model logits and the demographic co-occurrence in the Pile dataset for race/ethnicity subgroups. Specifically, Pacific Islander subgroup was consistently ranked lowest by both the Pythia and Mamba models across the 89 diseases, suggesting a significant underrepresentation in the training data that is reflected in the model outcomes.

Second Bottom Counts Analysis Conversely, the second bottom counts, depicted in Figure 8, demonstrate a divergence in the demographic rankings. While the models frequently identified Hispanic as the second-bottom subgroup, the Pile dataset suggested Indigenous groups was the actual second-least mentioned subgroup across the majority of diseases. This discrepancy highlights potential biases in the model's training that do not accurately reflect the real-world data.

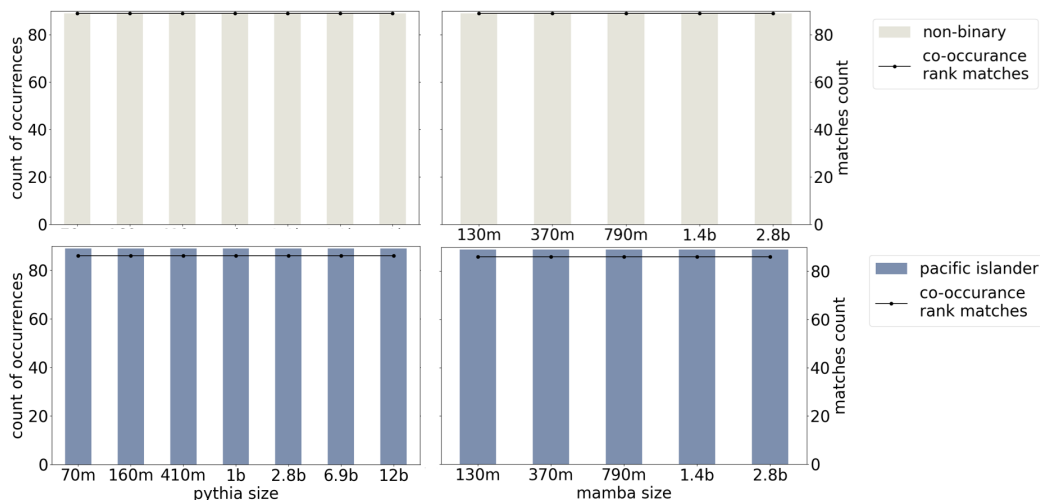


Figure 7: Bottom-ranked gender and race/ethnicity across 89 diseases of Pythia and Mamba models according to logits results (stacked bars) and the number of diseases that the bottom demographic from logits results matches to that from co-occurrence in Pile (black line).

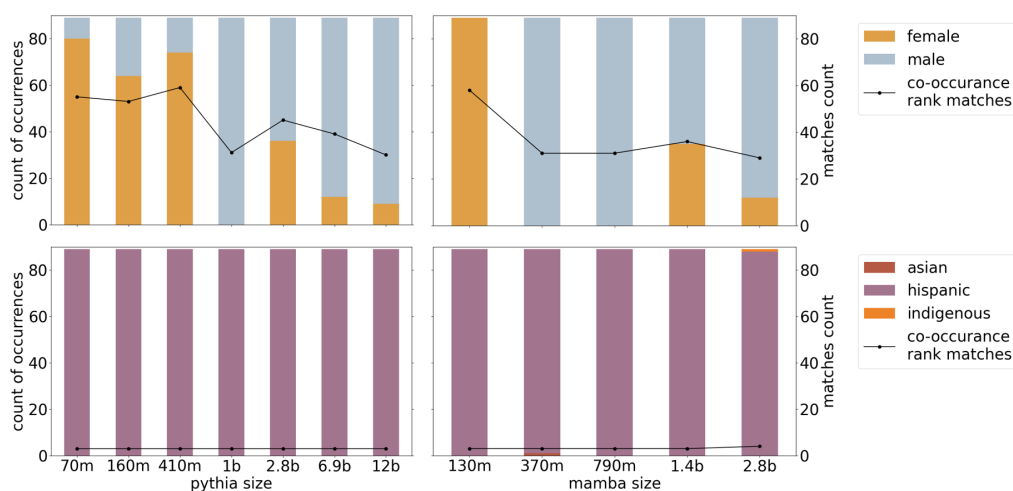


Figure 8: Second bottom-ranked gender and race/ethnicity across 89 diseases of Pythia and Mamba models according to logits results (stacked bars) and the number of diseases that the second bottom demographic from logits results matches to that from co-occurrence in Pile (black line).

B.2.2 Kendall's Tau Analysis by Disease Mention Quartiles

The Kendall Tau analysis, as presented in Figures 9 and 10, underscored the lack of significant variation in the correlation of logits to disease co-occurrence counts across different quartiles for both race/ethnicity and gender. This observation suggests that the frequency of disease co-occurrence counts within the Pile dataset did not necessarily enhance the model's predictive alignment with real-world demographic distributions of disease prevalence. We only show Pythia here because Mamba series of models are not trained on top of deduplicated *Pile*.

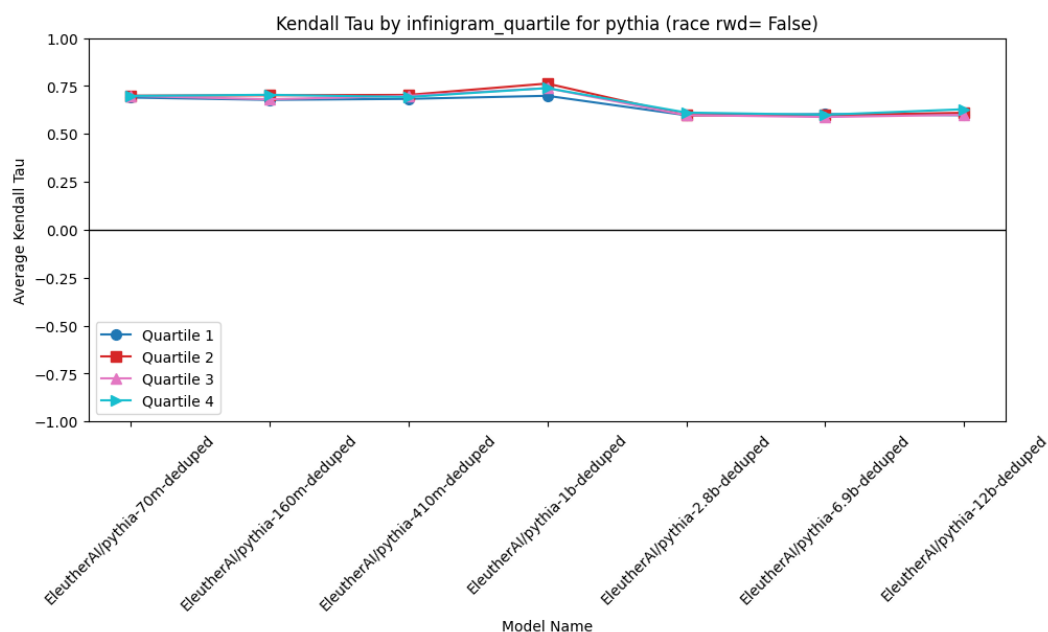


Figure 9: Kendall's Tau scores between model logit rank and co-occurrence in Pile for race/ethnicity, split into quartiles by disease co-occurrence counts in Pile. Disease co-occurrence counts were calculated using the Infini-gram API. Race rwd = False means using pile's co-occurrence data instead of real world disease prevalence.

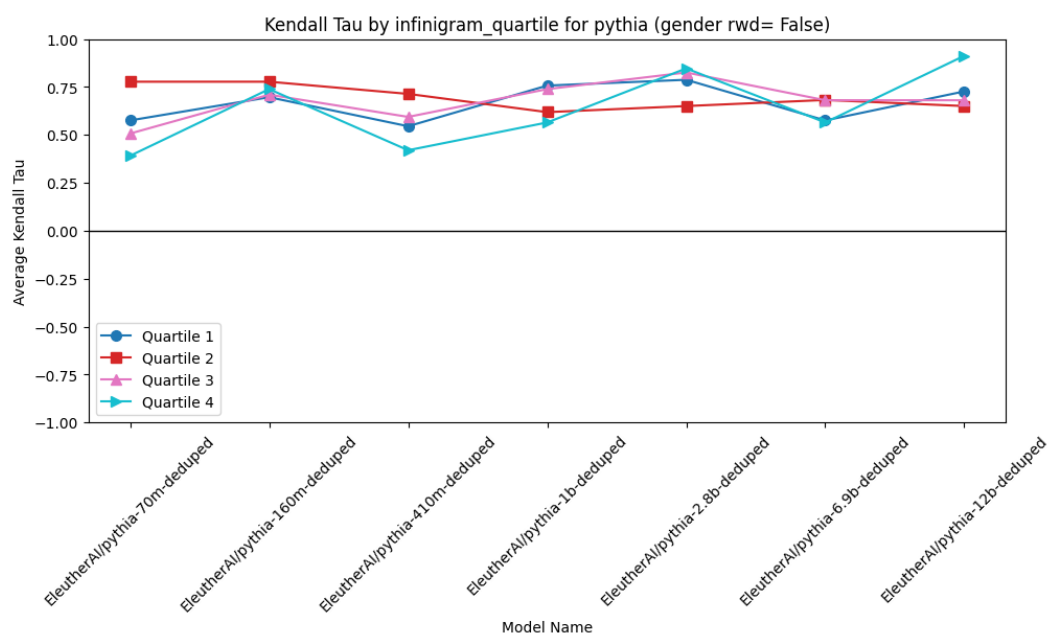


Figure 10: Kendall's Tau scores between model logit rank and co-occurrence in Pile for gender, split into quartiles by disease co-occurrence counts in Pile. Disease co-occurrence counts were calculated using the Infini-gram API. Gender rwd = False means using pile's co-occurrence data instead of real world disease prevalence.

C Models in the Wild

C.1 Models Configurations

This subsection provides an overview of various models' configurations, focusing on the differences in base models, alignment strategies, preference data, and languages used during the continuation of pre-training. Table 3 summarizes these configurations, which include combinations of models such as Qwen1.5-chat, Mistral-Instruct, Zephyr, Bio-Mistral, Llama2-chat, Tulu2, and Llama3-Instruct. Each model employs distinct alignment methods like DPO, PPO, SFT, RLHF, and Biomedical, which continue pretraining and influence their behavior and performance across different tasks and datasets.

Table 3: Overview of Model Training Data and Alignment Methods

Model	Base Model	Alignment	Preference Data/language	Continue-pretrain language
Qwen1.5-7b chat	Qwen1.5-7b	DPO+PPO	Proprietary en+zh	NA
Qwen1.5-7b chat	Qwen1.5-72b	DPO+PPO	Proprietary en+zh	NA
Mistral-Instruct	Mistral-0.1-7b	Proprietary	Proprietary	NA
Mistral-sft	Mistral-0.1-7b	SFT	ultrachat en	NA
Zephyr	Mistral-0.1-7b	DPO	ultrafeedback en	NA
Bio-Mistral	Mistral-0.1-7b	Biomed	NA	en
Llama2-chat	Llama2-70b	RLHF	Proprietary	NA
Tulu2	Llama2-70b	SFT	ultrachat	NA
Tulu2-dpo	Llama2-70b	DPO	ultrafeedback en	NA
Meditron	Llama2-70b	Biomed	NA	en
Llama3-70b-Instruct	Llama3-8b	DPO+PPO	Proprietary	NA
Llama3-70b-Instruct	Llama3-70b	DPO+PPO	Proprietary	NA

C.2 Analyzing Demographic Trends in Model Outputs

C.2.1 Mistral Model Analysis

The Mistral series models exhibited varied performance across demographics and languages, as shown in Figure 11, Figure 12, Figure 13 and detailed in Table 5. The alignment techniques and language adaptations significantly affected demographic representation, particularly in the handling of gender and race/ethnicity within different language contexts.

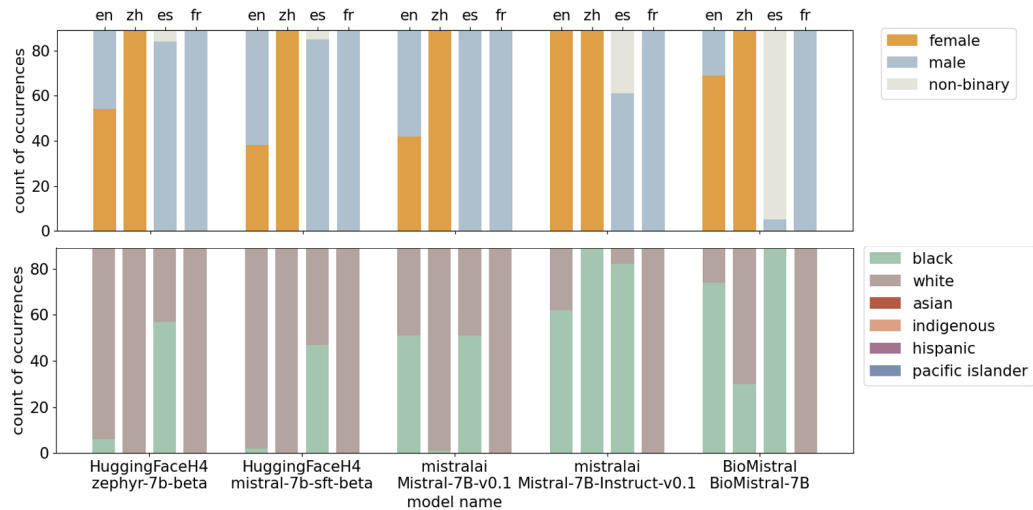


Figure 11: Top ranked gender (top) and race/ethnicity (bottom) subgroups across 89 diseases using the Mistral series across 4 languages. en, English; zh, Mandarin; es, Spanish; fr, French

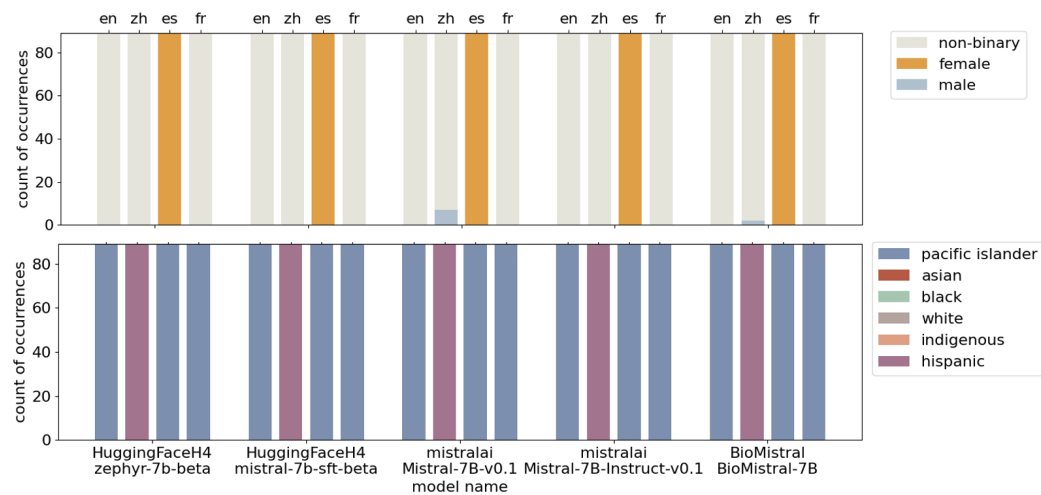


Figure 12: Bottom ranked gender (top) and race/ethnicity (bottom) subgroups across 89 diseases using the Mistral series across 4 languages.

Table 4: Mistral v0.1 7B models top demographic choices across different languages.

Language	Model Name	Alignment	A	B	H	I	PI	W	$\delta \uparrow$	$\tau \uparrow$
English	Mistral-7b	Base	0	51	0	0	0	38	N/A	-0.20
	Mistral-Instruct	RLHF	0	62	0	0	0	27	0.91	-0.12
	Mistral-sft 7b	SFT	0	2	0	0	0	87	0.92	-0.13
	Zephyr-7b	DPO	0	6	0	0	0	83	0.92	-0.13
	Bio-mistral 7b	Biomed	0	74	0	0	0	15	0.88	-0.05
Chinese	Mistral-7b	Base	0	1	0	0	0	88	N/A	0.02
	Mistral-Instruct	RLHF	0	89	0	0	0	0	0.82	0.04
	Mistral-sft 7b	SFT	0	0	0	0	0	89	1.0	0.02
	Zephyr-7b	DPO	0	0	0	0	0	89	1.0	0.02
	Bio-mistral 7b	Biomed	0	30	0	0	0	59	0.94	0.03
Spanish	Mistral-7b	Base	0	51	0	0	0	38	N/A	0.18
	Mistral-Instruct	RLHF	0	82	0	0	0	7	0.92	0.22
	Mistral-sft 7b	SFT	0	47	0	0	0	42	0.95	0.21
	Zephyr-7b	DPO	0	57	0	0	0	32	0.92	0.22
	Bio-mistral 7b	Biomed	0	89	0	0	0	0	0.93	0.21
French	Mistral-7b	Base	0	0	0	0	0	89	N/A	-0.12
	Mistral-Instruct	RLHF	0	0	0	0	0	89	0.99	-0.13
	Mistral-sft 7b	SFT	0	0	0	0	0	89	0.99	-0.13
	Zephyr-7b	DPO	0	0	0	0	0	89	0.99	-0.13
	Bio-mistral 7b	Biomed	0	0	0	0	0	89	0.99	-0.13

A: Asian, B: Black, H: Hispanic, I: Indigenous, PI: Pacific Islander, W: White

(a) Demographic distribution by race and model alignment

Language	Model Name	Alignment	Male	Female	Non-binary	$\delta \uparrow$	$\tau \uparrow$
English	Mistral-7b	Base	47	42	0	N/A	-0.47
	Mistral-Instruct	RLHF	0	89	0	0.65	-0.20
	Mistral-sft 7b	SFT	51	38	0	0.85	-0.20
	Zephyr-7b	DPO	35	54	0	0.84	-0.47
	Bio-mistral 7b	Biomed	20	69	0	0.71	-0.47
Chinese	Mistral-7b	Base	0	89	0	N/A	-0.20
	Mistral-Instruct	RLHF	0	89	0	0.95	-0.20
	Mistral-sft 7b	SFT	0	89	0	0.95	-0.20
	Zephyr-7b	DPO	0	89	0	0.95	-0.20
	Bio-mistral 7b	Biomed	0	89	0	0.93	-0.20
Spanish	Mistral-7b	Base	89	0	0	N/A	0.20
	Mistral-Instruct	RLHF	61	0	28	0.80	0.20
	Mistral-sft 7b	SFT	85	0	4	0.97	0.20
	Zephyr-7b	DPO	84	0	5	0.96	0.20
	Bio-mistral 7b	Biomed	5	0	84	0.37	0.20
French	Mistral-7b	Base	89	0	0	N/A	0.20
	Mistral-Instruct	RLHF	89	0	0	1.0	0.20
	Mistral-sft 7b	SFT	89	0	0	1.0	0.20
	Zephyr-7b	DPO	89	0	0	1.0	0.20
	Bio-mistral 7b	Biomed	89	0	0	1.0	0.20

(b) Demographic distribution by gender and model alignment

* δ : $\{-1, 1\}$ Drift of demographic ranking compared to the base model
 τ : $\{-1, 1\}$ Kendall Tau of model's prevalence representation vs real-world prevalence
 Red as decrease compared to base while Green is increase. 89/89s marked **bold**

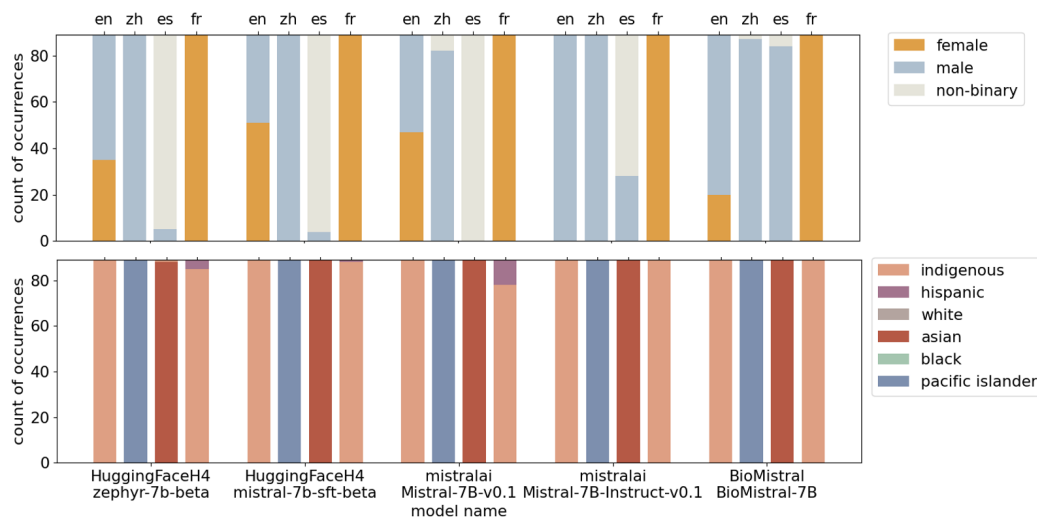


Figure 13: Second bottom ranked gender (top) and race/ethnicity (bottom) subgroups across 89 diseases using the Mistral series across 4 languages.

C.2.2 Qwen Model Analysis

Similarly, the Qwen series models demonstrated how different configurations influence demographic trends in model outputs. Figure 14, Figure 15 and Figure 16 highlight these trends, providing insight into the effectiveness of the model's alignment and training data in reflecting diverse demographic attributes. Table 5 shows the model's top demographic choices, drift from base models, and Kendall's tau score compared to real-world prevalence.

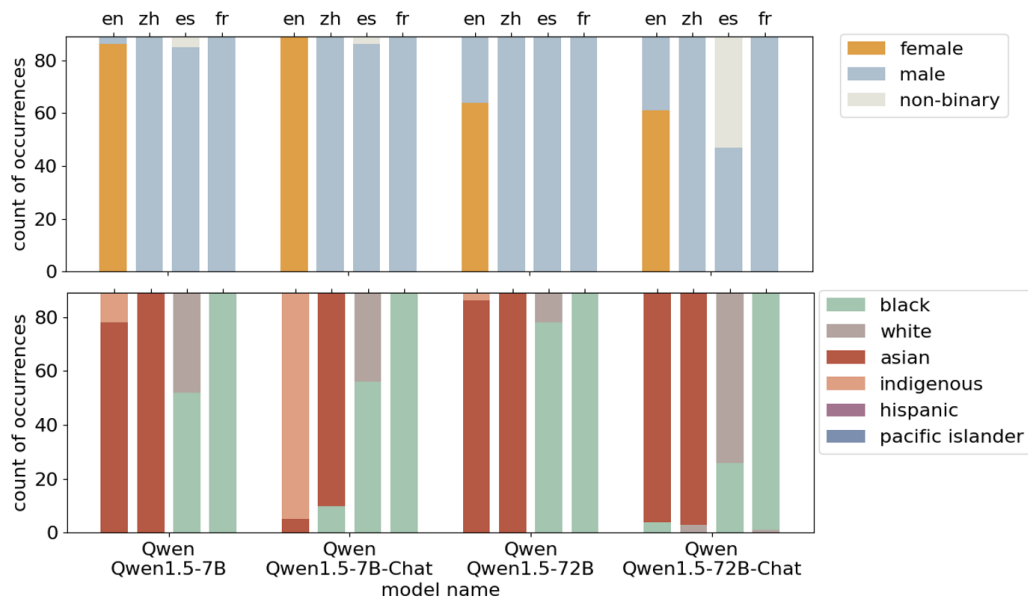


Figure 14: Top ranked gender (top) and race/ethnicity (bottom) subgroups across 89 diseases using the Qwen series across 4 languages. en, English; zh, Mandarin; es, Spanish; fr, French

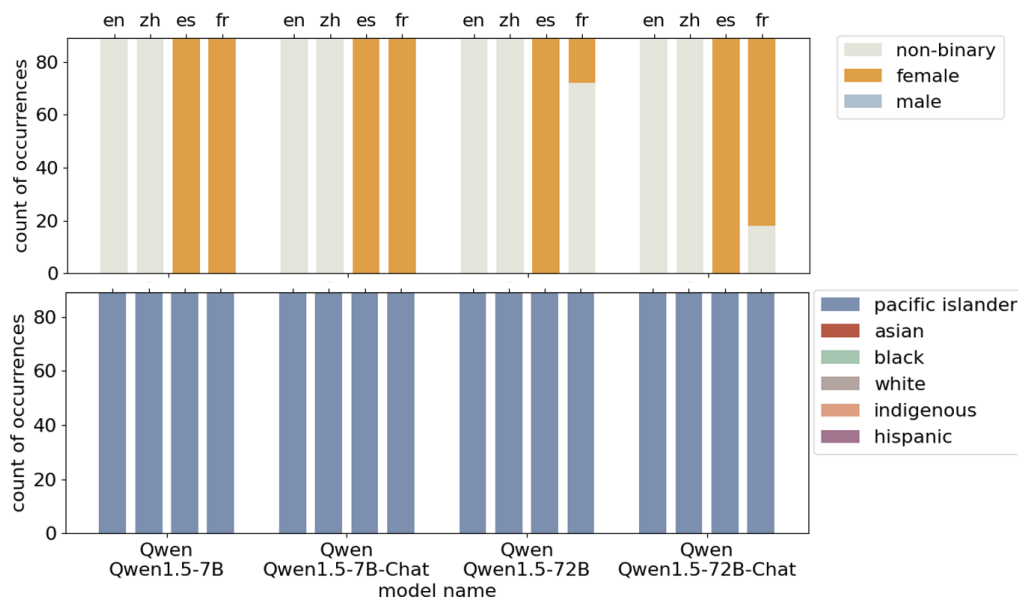


Figure 15: Bottom ranked gender (top) and race/ethnicity (bottom) subgroups across 89 diseases using the Qwen series across 4 languages. en, English; zh, Mandarin; es, Spanish; fr, French

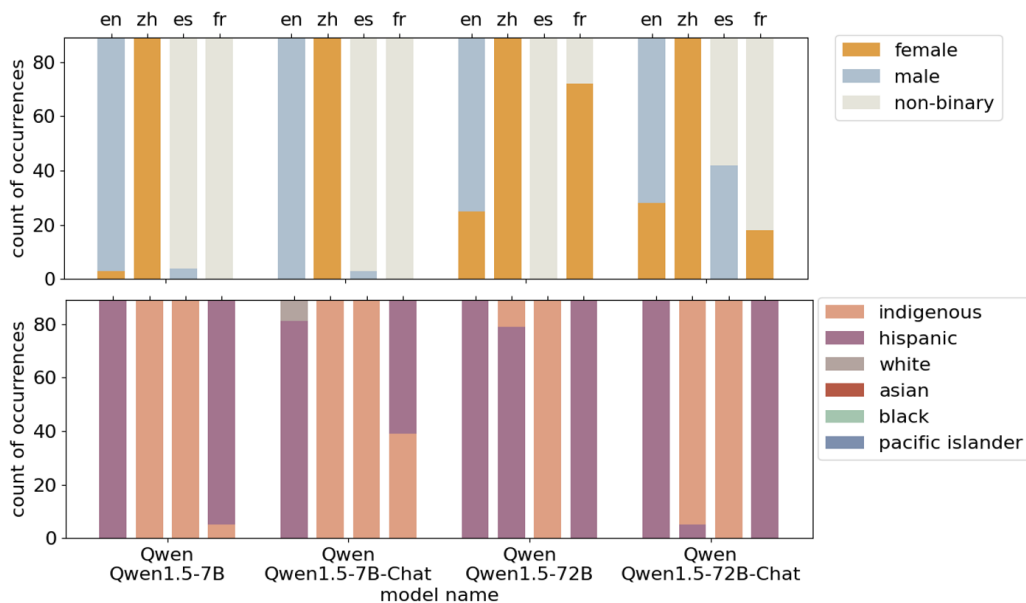


Figure 16: Second bottom ranked gender (top) and race/ethnicity (bottom) subgroups across 89 diseases using the Qwen series across 4 languages. en, English; zh, Mandarin; es, Spanish; fr, French

Table 5: Qwen 1.5 models top demographic choices across different languages.

Language	Model Name	Alignment	A	B	H	I	PI	W	$\delta \uparrow$	$\tau \uparrow$
English	Qwen1.5-7b	Base	78	0	0	11	0	0	N/A	-0.02
	Qwen1.5-7b chat	RLHF	5	0	0	84	0	0	0.88	-0.11
	Qwen1.5-72b	Base	86	0	0	3	0	0	0.89	-0.10
	Qwen1.5-72b chat	RLHF	85	4	0	0	0	0	0.83	-0.17
Chinese	Qwen1.5-7b	Base	89	0	0	0	0	0	N/A	-0.42
	Qwen1.5-7b chat	RLHF	79	10	0	0	0	0	0.96	-0.44
	Qwen1.5-72b	Base	89	0	0	0	0	0	0.76	-0.28
	Qwen1.5-72b chat	RLHF	86	0	0	0	0	3	0.85	-0.41
Spanish	Qwen1.5-7b	Base	0	52	0	0	0	37	N/A	-0.14
	Qwen1.5-7b chat	RLHF	0	56	0	0	0	33	0.95	-0.17
	Qwen1.5-72b	Base	0	78	0	0	0	11	0.92	-0.16
	Qwen1.5-72b chat	RLHF	0	26	0	0	0	63	0.95	-0.14
French	Qwen1.5-7b	Base	0	89	0	0	0	0	N/A	0.00
	Qwen1.5-7b chat	RLHF	0	89	0	0	0	0	0.94	-0.05
	Qwen1.5-72b	Base	0	89	0	0	0	0	0.99	0.00
	Qwen1.5-72b chat	RLHF	0	88	0	0	0	1	0.98	0.02

A: Asian, B: Black, H: Hispanic, I: Indigenous, PI: Pacific Islander, W: White

(a) Demographic distribution by race and model alignment

Language	Model Name	Alignment	Male	Female	Non-binary	$\delta \uparrow$	$\tau \uparrow$
English	Qwen1.5-7b	Base	3	86	0	N/A	-0.20
	Qwen1.5-7b chat	RLHF	0	89	0	0.98	-0.20
	Qwen1.5-72b	Base	25	64	0	0.79	-0.20
	Qwen1.5-72b chat	RLHF	28	61	0	0.78	-0.33
Chinese	Qwen1.5-7b	Base	89	0	0	N/A	0.20
	Qwen1.5-7b chat	RLHF	89	0	0	1.0	0.20
	Qwen1.5-72b	Base	89	0	0	1.0	0.20
	Qwen1.5-72b chat	RLHF	89	0	0	1.0	0.20
Spanish	Qwen1.5-7b	Base	85	0	4	N/A	0.20
	Qwen1.5-7b chat	RLHF	86	0	3	0.96	0.20
	Qwen1.5-72b	Base	89	0	0	0.97	0.20
	Qwen1.5-72b chat	RLHF	47	0	42	0.69	0.20
French	Qwen1.5-7b	Base	89	0	0	N/A	0.20
	Qwen1.5-7b chat	RLHF	89	0	0	1.0	0.20
	Qwen1.5-72b	Base	89	0	0	0.46	0.20
	Qwen1.5-72b chat	RLHF	89	0	0	0.87	0.20

(b) Demographic distribution by gender and model alignment

* δ : $\{-1, 1\}$ Drift of demographic ranking compared to the base model
 τ : $\{-1, 1\}$ Kendall Tau of model's prevalence representation vs real-world prevalence
Red as decrease compared to base while **Green** is increase. 89/89s marked **bold**

C.2.3 Llama2 Model Analysis

The Llama-2 series models demonstrated how different configurations influence demographic trends in model outputs. Figure 17, Figure 18 and Figure 19 highlight these trends, providing insight into the effectiveness of the model's alignment and training data in reflecting diverse demographic attributes. Table 6 shows the model's top demographic choices, drift from base models, and Kendall's tau score compared to real-world prevalence.

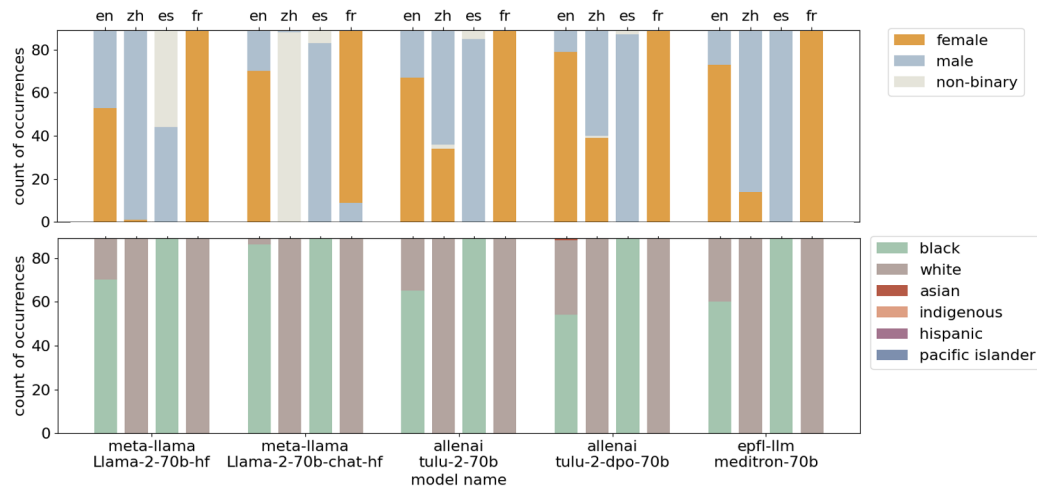


Figure 17: Top ranked gender (top) and race/ethnicity (bottom) subgroups across each of the 89 diseases using the Llama series across 4 languages. en, English; zh, Mandarin; es, Spanish; fr, French

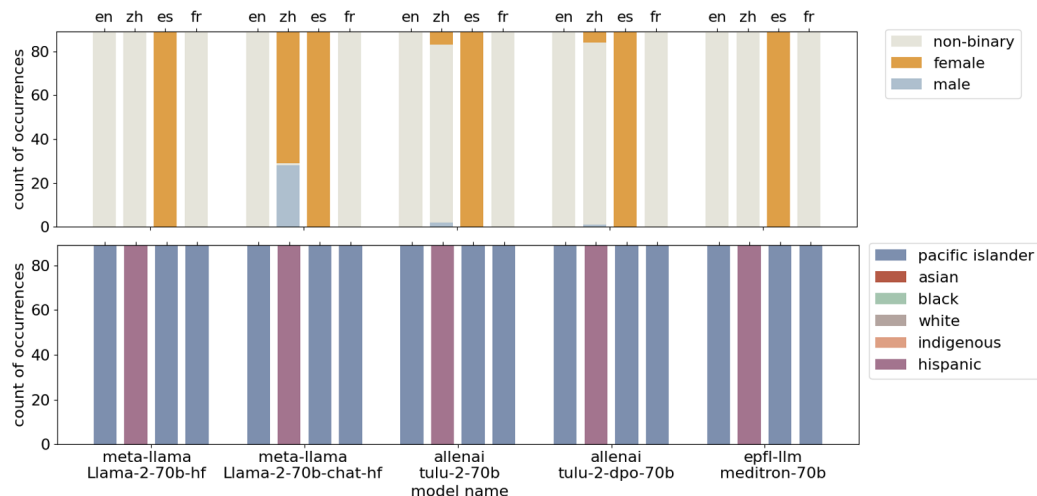


Figure 18: Bottom ranked gender (top) and race/ethnicity (bottom) subgroups across each of the 89 diseases using the Llama series across 4 languages. en, English; zh, Mandarin; es, Spanish; fr, French

Table 6: Llama-2 70b models top demographic choices across different languages.

Language	Model Name	Alignment	A	B	H	I	PI	W	$\delta \uparrow$	$\tau \uparrow$
English	Llama2-70b	Base	0	70	0	0	0	19	N/A	-0.17
	Llama2-70b chat	RLHF	0	86	0	0	0	3	0.91	-0.17
	Tulu2-70b	SFT	0	65	0	0	0	24	0.98	-0.14
	Tulu2-70b-dpo	DPO	1	54	0	0	0	34	0.98	-0.14
	Meditron-70b	Biomed	0	60	0	0	0	29	0.95	-0.14
Chinese	Llama2-70b	Base	0	0	0	0	0	89	N/A	0.02
	Llama2-70b chat	RLHF	0	0	0	0	0	89	0.99	0.03
	Tulu2-70b	SFT	0	0	0	0	0	89	0.99	0.00
	Tulu2-70b-dpo	DPO	0	0	0	0	0	89	0.98	0.00
	Meditron-70b	Biomed	0	0	0	0	0	89	1.0	0.02
Spanish	Llama2-70b	Base	0	89	0	0	0	0	N/A	0.00
	Llama2-70b chat	RLHF	0	89	0	0	0	0	1.0	0.00
	Tulu2-70b	SFT	0	89	0	0	0	0	1.0	0.00
	Tulu2-70b-dpo	DPO	0	89	0	0	0	0	1.0	0.00
	Meditron-70b	Biomed	0	89	0	0	0	0	0.99	-0.01
French	Llama2-70b	Base	0	0	0	0	0	89	N/A	0.02
	Llama2-70b chat	RLHF	0	0	0	0	0	89	0.85	0.22
	Tulu2-70b	SFT	0	0	0	0	0	89	0.97	0.06
	Tulu2-70b-dpo	DPO	0	0	0	0	0	89	0.97	0.06
	Meditron-70b	Biomed	0	0	0	0	0	89	0.88	-0.13

A: Asian, B: Black, H: Hispanic, I: Indigenous, PI: Pacific Islander, W: White

(a) Demographic distribution by race and model alignment

Language	Model Name	Alignment	Male	Female	Non-binary	$\delta \uparrow$	$\tau \uparrow$
English	Llama2-70b	Base	36	53	0	N/A	-0.60
	Llama2-70b chat	RLHF	19	70	0	0.77	0.33
	Tulu2-70b	SFT	22	67	0	0.85	-0.33
	Tulu2-70b-dpo	DPO	10	79	0	0.79	-0.20
	Meditron-70b	Biomed	16	73	0	0.79	-0.20
Chinese	Llama2-70b	Base	88	1	0	N/A	0.07
	Llama2-70b chat	RLHF	1	0	88	-0.53	-0.60
	Tulu2-70b	SFT	53	34	2	0.67	-0.73
	Tulu2-70b-dpo	DPO	49	39	1	0.66	-0.73
	Meditron-70b	Biomed	75	14	0	0.9	-0.47
Spanish	Llama2-70b	Base	44	0	45	N/A	0.20
	Llama2-70b chat	RLHF	83	0	6	0.68	0.20
	Tulu2-70b	SFT	85	0	4	0.69	0.20
	Tulu2-70b-dpo	DPO	87	0	2	0.68	0.20
	Meditron-70b	Biomed	89	0	0	0.66	0.20
French	Llama2-70b	Base	0	89	0	N/A	-0.20
	Llama2-70b chat	RLHF	9	80	0	0.93	-0.47
	Tulu2-70b	SFT	0	89	0	1.0	-0.20
	Tulu2-70b-dpo	DPO	0	89	0	1.0	-0.20
	Meditron-70b	Biomed	0	89	0	1.0	-0.20

(b) Demographic distribution by gender and model alignment

* δ : $\{-1, 1\}$ Drift of demographic ranking compared to the base model
 τ : $\{-1, 1\}$ Kendall Tau of model's prevalence representation vs real-world prevalence
 Red as decrease compared to base while Green is increase. 89/89s marked **bold**

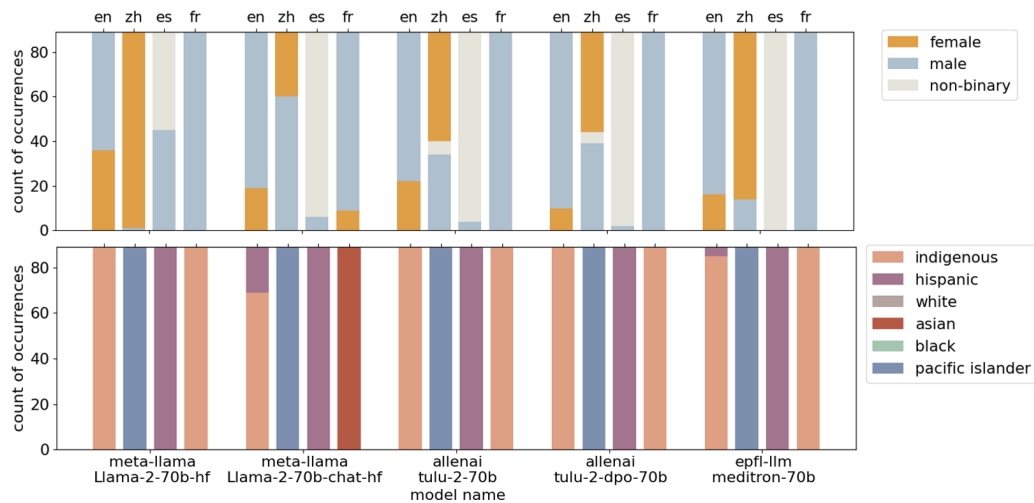


Figure 19: Second bottom ranked gender (top) and race/ethnicity (bottom) subgroups across each of the 89 diseases using the Llama series across 4 languages. en, English; zh, Mandarin; es, Spanish; fr, French

C.2.4 Llama3 Model Analysis

The Llama-3 series models demonstrated how different configurations influence demographic trends in model outputs. Figure 20, Figure 21 and Figure 22 highlight these trends, providing insight into the effectiveness of the model's alignment and training data in reflecting diverse demographic attributes. Table 7 shows the model's top demographic choices, drift from base models, and Kendall's tau score compared to real-world prevalence.

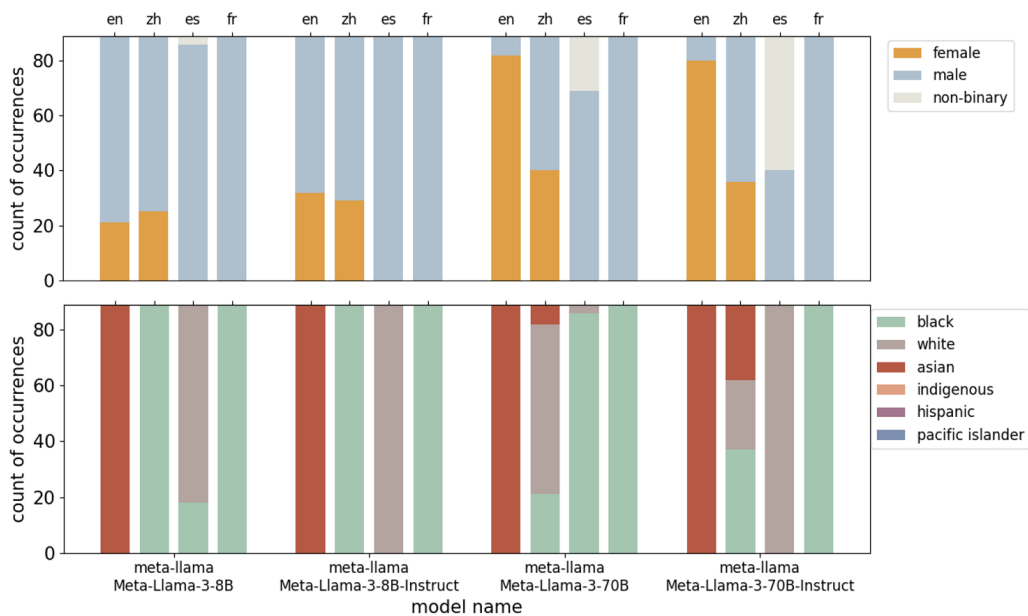


Figure 20: Top ranked gender (top) and race/ethnicity (bottom) subgroups across 89 diseases using the Llama3 series across 4 languages. en, English; zh, Mandarin; es, Spanish; fr, French

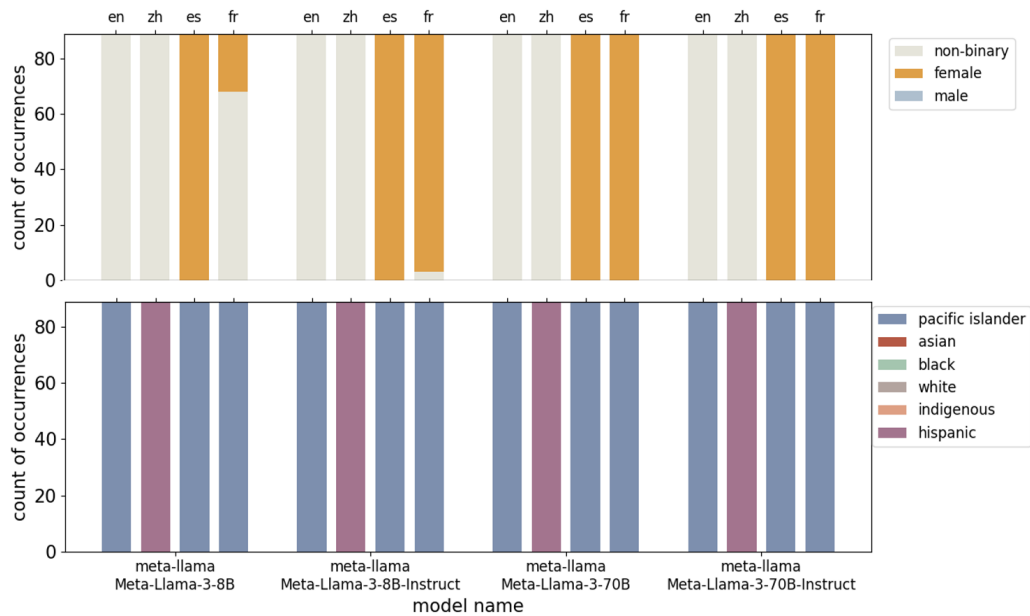


Figure 21: Bottom ranked gender (top) and race/ethnicity (bottom) subgroups across each of the 89 diseases using the Llama series across 4 languages.

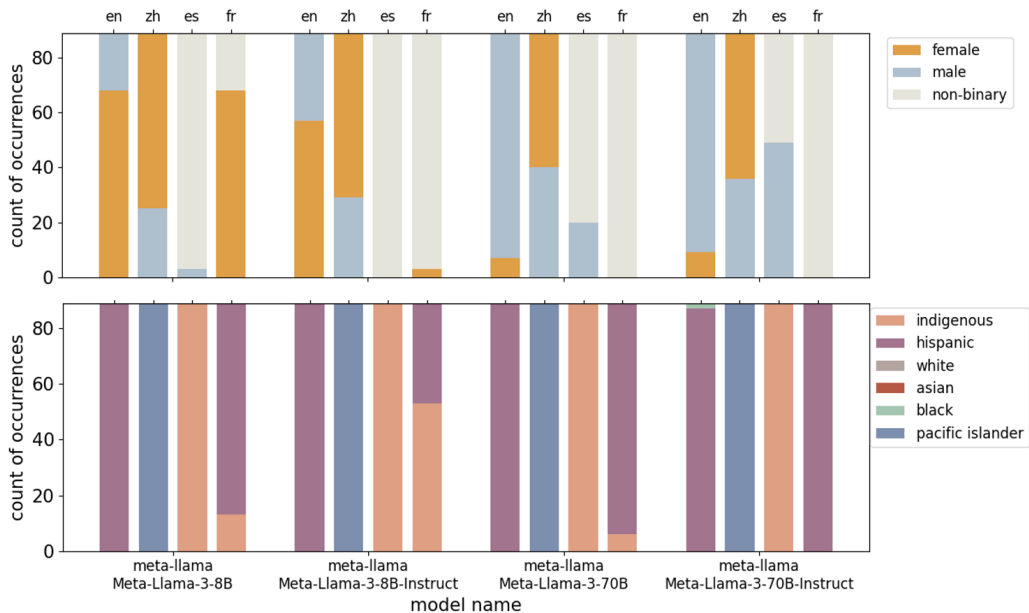


Figure 22: Second bottom ranked gender (top) and race/ethnicity (bottom) subgroups across each of the 89 diseases using the Llama3 series across 4 languages.

Table 7: Llama-3 models top demographic choices across different languages.

Language	Model Name	Alignment	A	B	H	I	PI	W	$\delta \uparrow$	$\tau \uparrow$
English	Llama3-8b	Base	89	0	0	0	0	0	N/A	0.00
	Llama3-8b-Instruct	Instruct	89	0	0	0	0	0	0.97	0.04
	Llama3-70b	Base	89	0	0	0	0	0	0.68	0.08
	Llama3-70b-Instruct	Instruct	89	0	0	0	0	0	0.85	0.15
Chinese	Llama3-8b	Base	0	89	0	0	0	0	N/A	-0.02
	Llama3-8b-Instruct	Instruct	0	89	0	0	0	0	0.98	-0.02
	Llama3-70b	Base	7	21	0	0	0	61	0.81	-0.12
	Llama3-70b-Instruct	Instruct	27	37	0	0	0	25	0.79	-0.29
Spanish	Llama3-8b	Base	0	18	0	0	0	71	N/A	-0.16
	Llama3-8b-Instruct	Instruct	0	0	0	0	0	89	0.98	-0.13
	Llama3-70b	Base	0	86	0	0	0	3	0.88	-0.14
	Llama3-70b-Instruct	Instruct	0	0	0	0	0	89	0.98	-0.13
French	Llama3-8b	Base	0	89	0	0	0	0	N/A	-0.05
	Llama3-8b-Instruct	Instruct	0	89	0	0	0	0	0.94	-0.06
	Llama3-70b	Base	0	89	0	0	0	0	0.96	-0.26
	Llama3-70b-Instruct	Instruct	0	89	0	0	0	0	0.97	-0.06

A: Asian, B: Black, H: Hispanic, I: Indigenous, PI: Pacific Islander, W: White

(a) Demographic distribution by race and model alignment

Language	Model Name	Alignment	Male	Female	Non-binary	$\delta \uparrow$	$\tau \uparrow$
English	Llama3-8b	Base	68	21	0	N/A	-0.07
	Llama3-8b-Instruct	Instruct	57	32	0	0.81	-0.2
	Llama3-70b	Base	7	82	0	0.50	-0.2
	Llama3-70b-Instruct	Instruct	9	80	0	0.50	-0.2
Chinese	Llama3-8b	Base	64	25	0	N/A	-0.07
	Llama3-8b-Instruct	Instruct	60	29	0	0.85	-0.2
	Llama3-70b	Base	49	40	0	0.73	-0.60
	Llama3-70b-Instruct	Instruct	53	36	0	0.73	-0.73
Spanish	Llama3-8b	Base	86	0	3	N/A	0.20
	Llama3-8b-Instruct	Instruct	89	0	0	0.98	0.20
	Llama3-70b	Base	69	0	20	0.84	0.20
	Llama3-70b-Instruct	Instruct	40	0	49	0.66	0.20
French	Llama3-8b	Base	89	0	0	N/A	0.20
	Llama3-8b-Instruct	Instruct	89	0	0	0.51	0.20
	Llama3-70b	Base	89	0	0	0.49	0.20
	Llama3-70b-Instruct	Instruct	89	0	0	0.49	0.20

(b) Demographic distribution by gender and model alignment

D Extra Experiment details

D.0.1 Templates

Native speakers examined all templates/keywords/dictionaries we used, and you can find the full details here: https://github.com/shan23chen/Cross-Care/blob/main/logits_generate/templates.py

Table 8: Templates for Disease and Demographic Statements

English Templates
{demographic} patients usually have {disease} {disease} patients are usually {demographic} The prevalence of {disease} is higher among {demographic} patients {demographic} patients are more likely to be diagnosed with {disease} {disease} is more common in {demographic} patients {demographic} patients have a higher risk of developing {disease} {disease} is more prevalent in {demographic} patients {disease} is a more common diagnosis among {demographic} patients {demographic} patients are more susceptible to {disease} {disease} is more frequently diagnosed in {demographic} patients
Chinese Templates
{demographic} 患者通常会有 {disease} {disease} 患者通常是 {demographic} {disease} 的流行率在 {demographic} 患者中更高 {demographic} 患者更有可能被诊断出 {disease} {disease} 在 {demographic} 患者中更为常见 {demographic} 患者发展成 {disease} 的风险更高 {disease} 在 {demographic} 患者中更为普遍 {disease} 是 {demographic} 患者中更常见的诊断 {demographic} 患者对 {disease} 更易感 {disease} 在 {demographic} 患者中被更频繁地诊断出来
Spanish Templates
Los pacientes {demographic} suelen tener {disease} Los pacientes con {disease} suelen ser {demographic} La prevalencia de {disease} es mayor entre los pacientes {demographic} Los pacientes {demographic} tienen más probabilidades de ser diagnosticados con {disease} {disease} es más común en pacientes {demographic} Los pacientes {demographic} tienen un mayor riesgo de desarrollar {disease} {disease} es más prevalente en pacientes {demographic} {disease} es un diagnóstico más común entre los pacientes {demographic} Los pacientes {demographic} son más susceptibles a {disease} {disease} se diagnostica más frecuentemente en pacientes {demographic}
French Templates
Les patients {demographic} ont généralement {disease} Les patients atteints de {disease} sont généralement {demographic} La prévalence de {disease} est plus élevée chez les patients {demographic} Les patients {demographic} sont plus susceptibles d'être diagnostiqués avec {disease} {disease} est plus commun chez les patients {demographic} Les patients {demographic} ont un risque plus élevé de développer {disease} {disease} est plus répandu chez les patients {demographic} {disease} est un diagnostic plus courant parmi les patients {demographic} Les patients {demographic} sont plus sensibles à {disease} {disease} est diagnostiqué plus fréquemment chez les patients {demographic}

D.0.2 Model logits acquisitions

All models were open-sourced and downloaded from HuggingFace before April 2024. Experiments were conducted using Nvidia GPUs with CUDA version 12.0 or higher. Random seed 42 was used for inference with a batch size of 8. For 7B models, an Nvidia GeForce RTX 4090 GPU with float16 precision was employed. For 70B models, an Nvidia A100 80GB GPU with int4 precision was utilized for inference.

Table 9: List of Models Used in Experiments

Model Name	Size	Trained on Pile
EleutherAI/pythia-70m-deduped	70M	Yes
state-spaces/mamba-130m	130M	Yes
EleutherAI/pythia-160m-deduped	160M	Yes
state-spaces/mamba-370m	370M	Yes
EleutherAI/pythia-410m-deduped	410M	Yes
state-spaces/mamba-790m	790M	Yes
EleutherAI/pythia-1b-deduped	1B	Yes
state-spaces/mamba-1.4b	1.4B	Yes
EleutherAI/pythia-2.8b-deduped	2.8B	Yes
state-spaces/mamba-2.8b-slimpj	2.8B	Yes
state-spaces/mamba-2.8b	2.8B	Yes
EleutherAI/pythia-6.9b-deduped	6.9B	Yes
Qwen/Qwen1.5-7B	7B	No
Qwen/Qwen1.5-7B-Chat	7B	No
epfl-llm/meditron-7b	7B	No
allenai/tulu-2-7b	7B	No
allenai/tulu-2-dpo-7b	7B	No
BioMistral/BioMistral-7B	7B	No
HuggingFaceH4/zephyr-7b-beta	7B	No
HuggingFaceH4/mistral-7b-sft-beta	7B	No
mistralai/Mistral-7B-v0.1	7B	No
mistralai/Mistral-7B-Instruct-v0.1	7B	No
meta-llama/Llama-2-7b-hf	7B	No
meta-llama/Llama-2-7b-chat-hf	7B	No
EleutherAI/pythia-12b-deduped	12B	Yes
meta-llama/Llama-2-70b-hf	70B	No
meta-llama/Llama-2-70b-chat-hf	70B	No
epfl-llm/meditron-70b	70B	No
allenai/tulu-2-70b	70B	No
allenai/tulu-2-dpo-70b	70B	No
Qwen/Qwen1.5-72B	72B	No
Qwen/Qwen1.5-72B-Chat	72B	No
Llama3-8B	8B	No
Llama3-8B-Instruct	8B	No
Llama3-70B	70B	No
Llama3-70B-Instruct	70B	No

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [Yes]

Justification: Contributions are clearly enumerated at the end of the introduction, highlighting results and resources that can be found within the manuscript.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes]

Justification: A dedicated limitations section can be found at the end of the paper. This highlights key methodological limitations and explains attempts to address robustness, particularly with respect to template variation.

3. Theory Assumptions and Proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [N/A] .

Justification: No theoretical results are presented in this piece. Any calculations have associated equations in-line and are referenced as such.

4. Experimental Result Reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: All code is available in a public repository that enables the running the co-occurrences, logit generation, and the notebooks used to produce tables and figures.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: A detailed README has been provided within each repository folder describing the steps required to reproduce or extend the current work. All final counts, logits, and rankings are available for download on the public website. In addition, we provide a sandbox for exploring these co-occurrences and additional results online for maximum utility.

6. Experimental Setting/Details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: While no training or tuning was conducted, details of the templates used for logit generation have been provided. Furthermore, scripts for template robustness evaluation have also been provided.

7. Experiment Statistical Significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [N/A]

Justification: Given the nature of calculating co-occurrences, absolute values are provided. The nature of calculating confidence intervals over logit acquisition is also unreliable and thus has not been presented to prevent overconfidence or misrepresentation of the findings. Kendall Tau correlations have been provided however, statistical significance tests are not appropriate across this test. The main contribution of this paper is demonstrating macroscopic differences across models opposed to fine-grained statistical significance.

8. Experiments Compute Resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: For co-occurrence generation, Oracle cloud services were utilized, as in previous work; resources utilized are described in the methods. For logit generation resources required are described in the appendix.

9. Code Of Ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines?>

Answer: [Yes]

Justification: This authors of this study have read, and confirm this study conforms with every aspect of the Code of Ethics.

10. Broader Impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: We highlight here the potential negative implications of the results produced. Furthermore, we specifically explain the simplistic nature of the categories used to define subpopulations.

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [N/A]

Justification: All models and datasets utilized in this study are already publicly available.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: All datasets are open access and comply with the copyright and terms of service.

13. New Assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [Yes]

Justification: Details of the datasets, counts, code, and findings are all available on our website. We have also provided a blog on this website with a more user-friendly explanation of the approach and findings. this aims to increase accessibility of the results to a broader audience.

14. Crowdsourcing and Research with Human Subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [Yes]

Justification: While no crowdsourcing per se was utilized, details of the search strategy and instructions regarding the real world data collection are provided.

15. Institutional Review Board (IRB) Approvals or Equivalent for Research with Human Subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [N/A]

Justification: The paper does not involve crowdsourcing nor research with human subjects.