

Performance in a dialectal profiling task of LLMs for varieties of Brazilian Portuguese

Raquel M. Ko Freitag¹², Túlio Sousa de Gois¹³

¹Laboratório Multiusuário de Informática e Documentação Linguística
Universidade Federal de Sergipe – Brazil (UFS)
Didática II – 49.107-230 – São Cristóvão – SE – Brazil

²Departamento de Letras Vernáculas – Universidade Federal de Sergipe – UFS

³Departamento de Computação – Universidade Federal de Sergipe – UFS

{rkofreitag,tuliosg}@academico.ufs.br

Abstract. *Different of biases are reproduced in LLM-generated responses, including dialectal biases. A study based on prompt engineering was carried out to uncover how LLMs discriminate varieties of Brazilian Portuguese, specifically if sociolinguistic rules are taken into account in four LLMs: GPT 3.5, GPT-4o, Gemini, and Sabiá-2. The results offer sociolinguistic contributions for an equity fluent NLP technology.*

Resumo. *Vieses de diferentes tipos são reproduzidos em respostas geradas por LLMs, inclusive dialetais. Um estudo baseado em engenharia de prompt foi realizado para descobrir como os LLMs discriminam as variedades do português brasileiro, especificamente se regras sociolinguísticas são consideradas por quatro LLMs – GPT 3.5, GPT-4o, Gemini e Sabiá-2 – na geração de suas respostas. Os resultados oferecem contribuições sociolinguísticas para uma tecnologia de PLN com equidade dialetal.*

1. Introduction

Advances in generative AI have enabled near-human responses, crucial for overcoming the Turing test [Danziger 2018]. However, achieving this requires algorithms to replicate ethically questionable human behaviors, including biases learned by large language models (LLMs) [Freitag 2021].

Biases can be explicit, consciously manipulated, or implicit, operating unconsciously through automatic associations. These biases affect generative AI in two key areas: the rules and filters applied during LLM fine-tuning, and the linguistic datasets used for training. However, the specifics of these biases—whether in rules, filters, or dataset selection—remain unclear [Bender et al. 2021]. To investigate these biases, reverse-engineering through prompt engineering is necessary, similar to how sociolinguistics studies human language attitudes.

In Brazil, sociolinguistic studies over the past 50 years have highlighted significant asymmetries between prestigious and non-standard varieties (whether regional or socially stigmatized), often perpetuated by implicit biases in educational materials and media, such as the portrayal of regional accents in soap operas [Freitag 2016]. These societal biases likely extend to LLMs.

For AI to be ethically and socially sensitive, the diversity of societal communities must be reflected in the language samples used to train LLMs. [Grieve et al. 2024] define a language variety as “a population of texts defined by external factors, such as the social background of the people who produce these texts, the social context in which these texts are produced, and the time period over which these texts are produced.”

Currently, there is no transparency on how language samples are collected and balanced to reflect linguistic diversity. A study using prompt engineering could reveal how LLMs handle varieties of Brazilian Portuguese and whether they consider sociolinguistic rules.

2. Dialectal biases in LLMs

Brazil’s continental size contributes to its dialectal diversity, further enriched by social diversity in language use. Both geolinguistic and sociolinguistic approaches have systematically described these patterns in Brazilian Portuguese [Roncarati et al. 2003, Abraçado and Martins 2015].

Though linguistic diversity doesn’t align strictly with socio-political boundaries, it is socially perceived and manifests in stereotypes, such as the classic *biscoito* vs. *bolacha* ‘cookie’ debate,¹ regional jokes and memes,² or even humorous maps that reflect aspects of perceptual dialectology [Preston 2010, Freitag et al. 2015, Freitag et al. 2016].

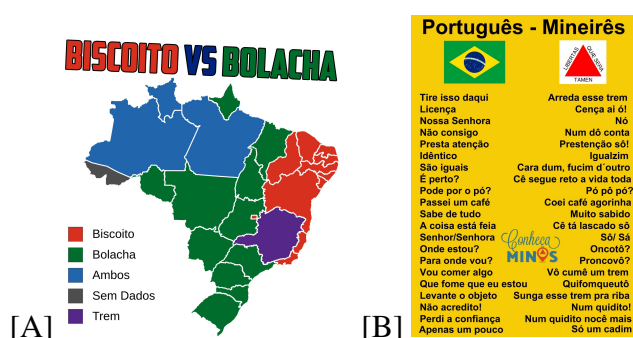


Figure 1. *Biscoito* or *bolacha* [A] and mineiros’s memes [B]

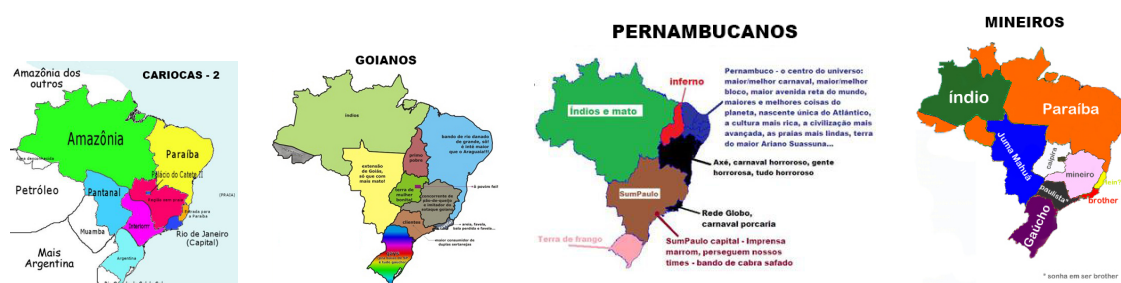


Figure 2. Funny maps with Brazilian stereotypes

Sociolinguistic surveys support these representations, showing systematic patterns between dialectal groups and territories, such as variations in first-person plural (*nós* a

¹<https://rionoticias.com.br/afinal-bolacha-ou-biscoito/>

²<https://www.conhecaminas.com/2016/10/10-coisas-que-so-mineiro-fala-entende-e.html>

gente) and second-person pronouns (*tu* *você*), their agreement (*nós vamos* *nós vai* — *tu vais* *tu vai*), and the shift between subjunctive and indicative moods for imperative forms (*cante* *canta*) [Abraçado and Martins 2015]. Phonological variations, like diminutives and diphthongs, also appear in written representations, particularly in memes.

It is expected that LLMs, like humans, will learn and reproduce societal biases and stereotypes, including linguistic biases [Shrawgi et al. 2024, Fleisig et al. 2024], as observed in LLMs trained on African American English [Mengesha et al. 2021, Dacon and Tang 2021, Dacon et al. 2022, Dacon 2022]. This paper evaluates the accuracy of LLMs in identifying Brazilian Portuguese dialectal profiles, assessing both their agreement with human judgments and consistency across different models: GPT-4o and GPT 3.5 (OpenAI), Sabiá-2 (Maritaca AI), and Gemini (Google AI). None of these LLMs disclose the size of their training corpus or the sources of their texts.

3. Method

The procedures encompassed three stages: generation, classification, and data wrangling. The dataset and analysis scripts are available at <https://osf.io/un8cw/>.

3.1. Target-profiles generation

The initial step involved instructing LLMs to generate text passages aimed at creating typical linguistic profiles for each of Brazil's 27 states. The prompt used was:

“Write a Facebook biography for my T-shirt store in XXX language, replacing XXX with the name of the state.”

Data was collected on 10 June 2024 and then analyzed for linguistic dialectal features to identify dialectal clues.

3.2. Target-profiles classification

LLMs were asked to identify the state of the text generated in the first step (*task*) by two prompts:

- *task + input*
*“This is a target audience identification task. Looking at the linguistic marks, identify for target audience in which state of Brazil the INPUT advertisement was constructed:
INPUT: XXXX”*
- *task + clue* based on features identified by the human-judges + *input*
*“This is a target audience identification task. Looking at the linguistic clues, identify for the target audience in which state of Brazil the has been constructed:
INPUT: XXXX
CLUES: interjections, second person pronoun, first person plural pronoun, second person verb agreement, first person plural verb agreement ”*

3.3. Data Wrangling

Data from the classification tasks was submitted to three stages of processing: 1) **Data cleaning**: city names were removed, retaining only state names; 2) **Standardization**: null answers were standardized as “não identificado,” while errors like “Sudoeste” and “Região Amazônica” were retained; and 3) **Data structuring**: all classification data was categorized and stored in .csv files.

The models generated a total of 108 texts, 4 for each Brazilian state (1 per model) and 216 classifications were carried out, 108 for each type of prompt (*task* and *task + clue*). The generated texts were also classified by three human experts in syntactic dialectal variation, with the task of identifying the target state of the text. Measurement of the classifications (both of LLMs and humans) were carried out using Fleiss' Kappa [Fleiss 1971].

4. Results

4.1. Exploratory analysis of target-profiles

The exploratory analysis revealed that Sabiá-2 [Sales Almeida et al. 2024], a Brazilian Portuguese-trained LLM from Maritalk, showed no dialectal variability in its responses.

- Sabiá-2 generation example:

“Oi, [NOME DO LUGAR] Você já conhece a nossa loja de camisetas sustentáveis? Se ainda não, prepare-se para descobrir um novo jeito de expressar o seu amor pela nossa terra e pelo nosso planeta! Na nossa produção, usamos algodão orgânico, que respeita o ciclo da vida e a terra que pisa o nosso pé, e tintas naturais, que não sujam as águas que correm pelo nosso estado. As estampas são uma homenagem ao nosso estado: cada ponto turístico, cada paisagem natural, cada aspecto da nossa rica cultura é representado com orgulho e carinho. Você vai encontrar desde o estilo básico até o mais ousado, mas sempre com aquele toque eco-friendly que só nós sabemos dar. E aí, vai deixar de fazer parte dessa revolução verde? Na nossa loja, além de encontrar camisetas que fazem a diferença, você vai ter a certeza de que cada compra fortalece a nossa comunidade e a nossa luta por um mundo mais sustentável. Venha nos visitar, escolha a sua favorita e mostre que em [NOME DO LUGAR], além de um povo acolhedor e de uma natureza exuberante, temos a consciência ambiental como uma das nossas maiores riquezas! Com amor e orgulho [NOME DO LUGAR], [NOME DA LOJA]”

In contrast, GPT 3.5, GPT-4o, and Gemini displayed sensitivity to dialectal differences, evident in variations of 2P pronouns (*tu* *você*), 1PP agreement (*nós vai* *nós vamos*), subjunctive/indicative morphology (*compre* *compra*), interjections (*uai*, *oxente*, *ôxe*), and discourse markers (*visse?* *tá ligado?* *é nóix*, *bora lá*).

- Gemini generation example:

“E aí, pessoal! Bem-vindos à [NOME DA LOJA]! Aqui no coração da Amazônia, nós tá fazendo moda de um jeito sustentável e cheio de orgulho da nossa terra. Nossas camisas são feitas com materiais reciclados e tinturas naturais, tudo bem cuidadoso pra respeitar o meio ambiente. O negócio é ficar estiloso e cuidar da nossa floresta ao mesmo tempo, visse? Aqui você encontra camisa que é a cara do [NOME DO LUGAR]: cheia de cor, vida e consciência. Vem com a gente nessa jornada de vestir a sustentabilidade com muito orgulho e amor pela natureza. Bora lá, junte-se a nós e venha conferir as novidades!”

- GPT 3.5 generation example:

“E aí, meu chegado?! Seja bem-vindo à [NOME DA LOJA], a loja onde a moda sustentável é raiz! Aqui, o negócio é respeitar a floresta e o nosso mundão, com camisetas que são a cara do [NOME DO LUGAR]? Nossas camisas são tudo de bom, feitas com material reciclado e tingidas do jeito mais consciente, porque a gente se importa com a Amazônia e com a natureza que tá no nosso sangue, né não? Então, se tu quer se vestir com estilo e ainda ajudar a proteger o planeta, vem com a gente! Aqui tu vai encontrar o visual perfeito e ainda fazer bonito com o meio ambiente. Quando quiser, dá um pulo aqui na loja que a porta tá sempre aberta pra ti, meu amigo! Fica à vontade, viu?”

- GPT-4o generation example:

“E aí, galera! Bem-vindos à [NOME DA LOJA]! Aqui é onde o estilo se encontra com o respeito pela natureza, do jeitinho que [NOME DO LUGAR] merece! Nós trabalhamos com material reciclado e tingimento responsável, tudo feito com carinho pras nossas terras [NOME DO LUGAR] e pro nosso planetão. É um compromisso sério, viu? Então, vem conferir nossas camisas que são um arraso, cheias de estilo e consciência ecológica. É moda que valoriza o que é nosso e cuida do futuro, firmeza? Bora junto nessa ideia de mostrar que em [NOME DO LUGAR], o verde tá na moda!”

These features align with the descriptive patterns identified by previous sociolinguistic studies [Abraçado and Martins 2015], reinforcing that some LLMs learn linguistic biases.

The generated responses were evaluated by three human judges specializing in syntactic dialectal variation. They identified the target state and provided reasoning for each input. Although their assessments highlighted similar dialectal clues as the exploratory analysis, the agreement was weak (inter-annotator: $\kappa = 0.31$; target: $\kappa = 0.13$). The next step was to have the LLMs perform the same task.

4.2. Target-profiles evaluation

LLMs were tasked with identifying the state of the text generated in the first step using two prompts: 1) *task + input* and 2) *task + clue + input*. The analysis flow is shown in Figure 3 and displays the Brazilian states, the LLMs used in both generation and classification, and the classifications performed with the two types of prompt. Although all prompts were in Brazilian Portuguese, Gemini responded in English.

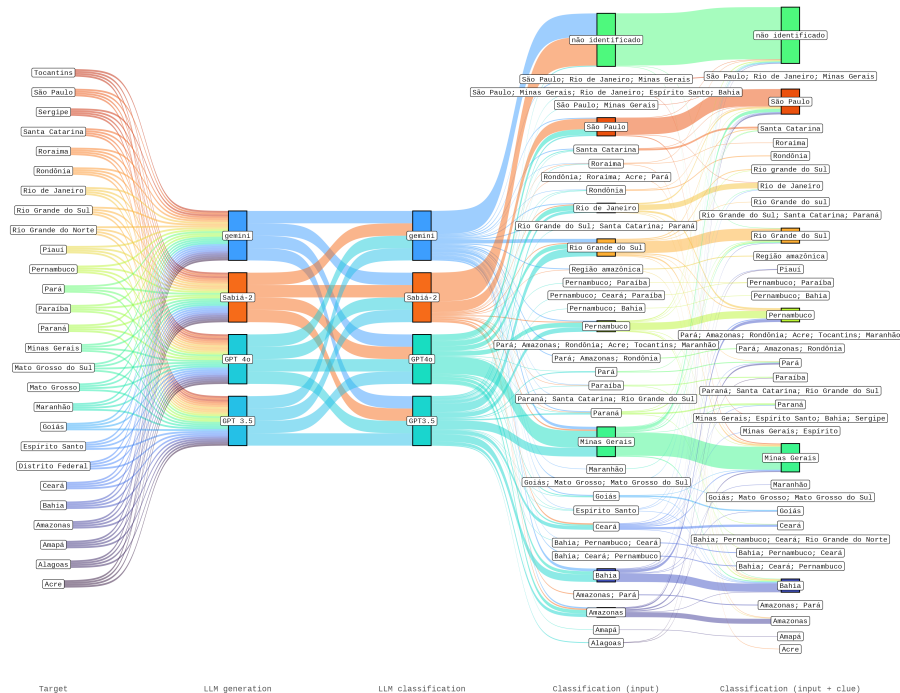


Figure 3. Analysis flowchart

In task 1, with weak agreement ($\kappa = 0.15$), differences emerged among LLMs: Sabiá-2 had the lowest agreement rate ($\kappa = 0.10$), while GPT-4o achieved the highest agreement, but still rather low ($\kappa = 0.21$). Identification rates varied as well: Sabiá-2 failed to identify a profile in 91% of cases, followed by Gemini at 88%, whereas GPT 3.5 and GPT-4o correctly identified the locations in over 20% of cases. Unexpectedly, in task 2, the non-identification rate increased (not significant), suggesting that the linguistic features identified by humans did not significantly influence the LLMs' classification. The only exception was GPT 3.5, which correctly identified one more location (an increase of 0.92%).

- This is an example of Sabiá-2's answer for task 1:

“Observando as marcas linguísticas e o conteúdo da propaganda, podemos determinar que o público-alvo é de um estado brasileiro com características culturais e ambientais específicas, que valoriza a sustentabilidade e tem um modo de falar próprio. [...] Considerando todas essas informações, é possível inferir que o público-alvo da propaganda é de um estado do Brasil com características culturais e ambientais específicas, que valoriza a sustentabilidade e tem um modo de falar próprio. Sem uma especificação clara do [NOME DA LOJA] ou do [NOME DO LUGAR], não é possível determinar exatamente qual estado é o público-alvo, mas as pistas sugerem um local com características semelhantes à região nordeste, dada a informalidade e o uso de gírias, ou talvez um estado do sul ou sudeste que tenha uma comunidade engajada em questões ambientais e use redes sociais como parte de sua estratégia de marketing.”

The input lacks specific geographic identification; it could be attributed to the “Nordeste” region due to informal style and slang, or to the “Sul” or “Sudeste” regions, presumed to have higher engagement in social issues. GPT 3.5, while imprecise in geographic localization, provides a more detailed identification of grammatical features even without additional clues:

“Com base nas marcas linguísticas presentes na propaganda do INPUT:

- **Uso de “aí” e “tá ligado”:** Essas expressões são comuns em diversas regiões do Brasil, mas especialmente no contexto da propaganda, elas podem indicar uma influência mais forte de dialetos urbanos, como os encontrados em estados do Sudeste e Centro-Oeste.
- **Realização do pronome de primeira pessoa do plural:** “Nóis trabalha firme...”, indicando um dialeto regional característico, comum em estados como São Paulo e Minas Gerais, por exemplo.
- **Expressões regionais como “mais da hora”:** Esta expressão é típica de regiões como o Sudeste e o Sul do Brasil.
- **Referência ao compromisso com o planeta:** A preocupação com sustentabilidade e meio ambiente é um tema geralmente valorizado em todo o Brasil, mas a forma como é apresentada na propaganda pode refletir uma sensibilidade urbana mais específica.

*Com base nessas características linguísticas, a propaganda do INPUT parece ser direcionada principalmente para um público-alvo urbano, possivelmente das regiões **Sudeste ou Centro-Oeste do Brasil**, onde há uma mistura de expressões informais e um foco crescente em sustentabilidade.”*

GPT-4o provides the most accurate identification of the state intended by the input. The subsequent answer relates to the same input as those given by GPT 3.5 and Sabiá-2:

“Observando as marcas linguísticas presentes na propaganda, podemos identificar o público-alvo provável com base nos seguintes elementos:

- **Interjeições e expressões coloquiais:** "E aí", "tá ligado?", "nóis trabalha", "da hora", "bora", "sim senhor", "tá plantado".
- **Realização do pronome de segunda pessoa:** "você" não é explicitamente utilizado, mas a forma de tratamento é informal e coloquial.
- **Realização do pronome de primeira pessoa do plural:** "nóis" em vez de "nós".
- **Concordância verbal com segunda pessoa:** Não é explicitamente mostrada.
- **Concordância verbal com primeira pessoa do plural:** "nóis trabalha" em vez de "nós trabalhamos".

Essas marcas linguísticas são típicas do estado de São Paulo, especialmente da capital e regiões metropolitanas. O uso de "tá ligado?", "da hora" e "nóis" são característicos do dialeto paulista, que muitas vezes é associado à cultura urbana paulistana. Portanto, a propaganda do INPUT parece ser construída para o público-alvo do estado de São Paulo."

GPT-4o, considering all the grammatical features, determines that the input is intended for São Paulo, providing a more precise answer than GPT 3.5. Like humans, GPT-4o integrates linguistic clues and covariation, as illustrated in the following example:

"A análise das marcas linguísticas presentes no texto da propaganda do INPUT sugere que o público-alvo é o estado de Pernambuco. Aqui estão as razões:

1. **Interjeições e expressões regionais:** A expressão "visse?" é tipicamente pernambucana e nordestina.
2. **Realização do pronome de segunda pessoa:** A forma "você" é usada, que é comum em várias regiões do Brasil.
3. **Realização do pronome de primeira pessoa do plural:** "Nóis" em vez de "nós" indica uma variação regional.
4. **Concordância verbal com segunda pessoa:** "deixar você", "ajudar nosso planeta".
5. **Concordância verbal com primeira pessoa do plural:** "Nóis trabalha", uma forma de falar mais informal e regional.

Outras expressões como "galera", "arretadas de boas" e "bora" reforçam ainda mais o tom informal e típico da linguagem usada no estado de Pernambuco. Portanto, a propaganda do INPUT foi construída para o público-alvo de Pernambuco, usando elementos linguísticos característicos da região para criar uma conexão mais forte com os consumidores locais."

The feature associated with urbanity in the previous answer is seen as a regional trait in this case, leading to the attribution of the input to Pernambuco. This variation can be explained by indexical fields [Eckert 2008]. For instance, *nós vai* is considered an informal feature for Pernambuco but an indicator of urban culture for São Paulo. This pattern aligns with societal studies of Brazilian Portuguese [Freitag et al. 2015].

In task 1), the most common states identified were Minas Gerais (9% for GPT-4o and 4% for GPT 3.5), São Paulo (5% for Sabiá-2), and Rio Grande do Sul (2% for Gemini) (Figure 4). In task 2), results were similar except GPT 3.5, which identified Pernambuco (Figure 5). Only GPT 3.5 and GPT-4o provided explicit analyses of clues, such as *tu* being common in Southern Brazil or *nóis vai* indicating informality or urban features. These findings highlight the sociolinguistic fine-tuning of LLMs or their *language regard* [Preston 2010].

5. Discussion

This study investigated dialectal sensitivity in LLMs by assessing their responses to tasks aimed at identifying dialectal features in Brazilian Portuguese.

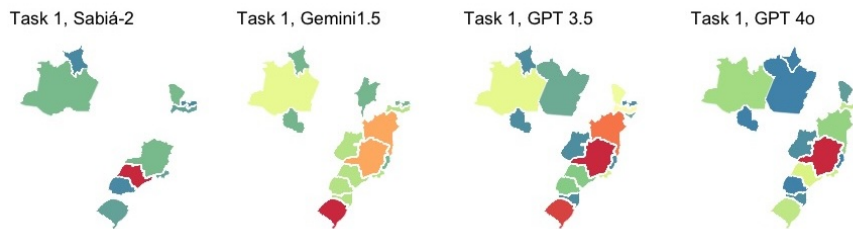


Figure 4. Geographical distribution of profile identification in task 1 *prompt: task + input*

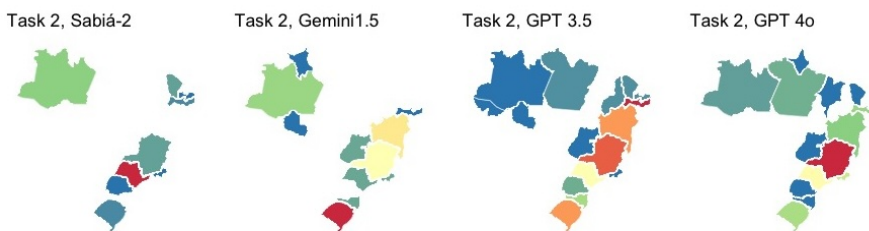


Figure 5. Geographical distribution of profile identification in task 2 *prompt: task + clue + input*

The findings revealed that Sabiá-2, a chatbot from Maritalk trained in Brazilian Portuguese, lacked dialectal sensitivity, showing no variability in responses. In contrast, GPT 3.5, GPT-4o, and Gemini demonstrated sensitivity to dialectal differences, evident in variations in pronoun usage, verb agreement, and other linguistic markers, aligning with sociolinguistic patterns. The agreement among human evaluators on dialectal features was weak ($\kappa = 0.31$), suggesting either inconsistent identification or a low number of judges. Among LLMs, Sabiá-2 had the lowest agreement rate ($\kappa = 0.08$), while GPT-4o showed the best agreement among the classifications ($\kappa = 0.21$). Notably, Gemini responded in English despite prompts being in Portuguese.

Incorporating specific linguistic features into prompts did not notably improve the LLMs' ability to identify the state, indicating that these features may not significantly affect classification. In task 1), GPT-4o and GPT 3.5 often identified Minas Gerais, while São Paulo and Rio Grande do Sul were identified by Sabiá-2 and Gemini1.5, respectively. In task 2), GPT 3.5 shifted its identification to Pernambuco. Only GPT 3.5 and GPT-4o provided explicit justifications based on dialectal clues, indicating some understanding of regional features.

These results show that while LLMs can detect dialectal variation, their ability to pinpoint specific regional profiles is inconsistent. The use of human-identified linguistic clues does not significantly enhance classification accuracy. Understanding how LLMs handle language varieties can help sociolinguistics explore human processing of linguistic variation and contribute to advancing linguistic justice and equitable NLP technologies [Baugh 2018, Wolfram and Eisenhauer 2019, Nee et al. 2021, Nee et al. 2022, Liu et al. 2023].

References

- Abraçado, J. and Martins, M. A. (2015). *Mapeamento sociolinguístico do português brasileiro*. Editora Contexto.
- Baugh, J. (2018). *Linguistics in pursuit of justice*. Cambridge University Press.
- Bender, E. M., Gebru, T., McMillan-Major, A., and Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? In *Proceedings of the 2021 ACM conference on fairness, accountability, and transparency*, pages 610–623.
- Dacon, J. (2022). Towards a deep multi-layered dialectal language analysis: A case study of african-american english. *arXiv preprint arXiv:2206.08978*.
- Dacon, J., Liu, H., and Tang, J. (2022). Evaluating and mitigating inherent linguistic bias of african american english through inference. In *Proceedings of the 29th International Conference on Computational Linguistics*, pages 1442–1454.
- Dacon, J. and Tang, J. (2021). What truly matters? using linguistic cues for analyzing the# blacklivesmatter movement and its counter protests: 2013 to 2020. *arXiv preprint arXiv:2109.12192*.
- Danziger, S. (2018). Where intelligence lies: Externalist and sociolinguistic perspectives on the turing test and ai. In *Philosophy and Theory of Artificial Intelligence 2017*, pages 158–174. Springer.
- Eckert, P. (2008). Variation and the indexical field. *Journal of sociolinguistics*, 12(4):453–476.
- Fleisig, E., Smith, G., Bossi, M., Rustagi, I., Yin, X., and Klein, D. (2024). Linguistic bias in chatgpt: Language models reinforce dialect discrimination. *arXiv preprint arXiv:2406.08818*.
- Fleiss, J. L. (1971). Measuring nominal scale agreement among many raters. *Psychological bulletin*, 76(5):378.
- Freitag, R. M. K. (2016). Sociolinguística no/do brasil. *Cadernos de Estudos Linguísticos*, 58(3):445–460.
- Freitag, R. M. K. (2021). Preconceito linguístico para humanizar as máquinas. *Cadernos de Linguística*, 2(4):e495–e495.
- Freitag, R. M. K., Severo, C. G., Rost-Snichelotto, C. A., and Tavares, M. A. (2015). Como o brasileiro acha que fala? desaios e propostas para a caracterização do” português brasileiro”. *Signo y seña*, (28):65–87.
- Freitag, R. M. K., Severo, C. G., Rost-Snichelotto, C. A., and Tavares, M. A. (2016). Como os brasileiros acham que falam? percepções sociolinguísticas de universitários do sul e do nordeste. *Todas as Letras-Revista de Língua e Literatura*, 18(2).
- Grieve, J., Bartl, S., Fuoli, M., Grafmiller, J., Huang, W., Jawerbaum, A., Murakami, A., Perlman, M., Roemling, D., and Winter, B. (2024). The sociolinguistic foundations of language modeling. *arXiv preprint arXiv:2407.09241*.
- Liu, Y., Held, W., and Yang, D. (2023). Dada: Dialect adaptation via dynamic aggregation of linguistic rules. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 13776–13793.

- Mengesha, Z., Heldreth, C., Lahav, M., Sublewski, J., and Tuennerman, E. (2021). “i don’t think these devices are very culturally sensitive.”—impact of automated speech recognition errors on african americans. *Frontiers in Artificial Intelligence*, 4:725911.
- Nee, J., Macfarlane Smith, G., Sheares, A., and Rustagi, I. (2021). Advancing social justice through linguistic justice: Strategies for building equity fluent nlp technology. In *Proceedings of the 1st ACM Conference on Equity and Access in Algorithms, Mechanisms, and Optimization*, pages 1–9.
- Nee, J., Smith, G. M., Sheares, A., and Rustagi, I. (2022). Linguistic justice as a framework for designing, developing, and managing natural language processing tools. *Big Data & Society*, 9(1):20539517221090930.
- Preston, D. R. (2010). Language, people, salience, space: perceptual dialectology and language regard. *Dialectologia: revista electrònica*, pages 87–131.
- Roncarati, C., Abraçado, J., and Heye, J. B. (2003). *Português brasileiro: contato lingüístico, heterogeneidade e história*, volume 2. 7 Letras.
- Sales Almeida, T., Abonizio, H., Nogueira, R., and Pires, R. (2024). Sabiá-2: A new generation of portuguese large language models. *arXiv e-prints*, pages arXiv–2403.
- Shrawgi, H., Rath, P., Singhal, T., and Dandapat, S. (2024). Uncovering stereotypes in large language models: A task complexity-based approach. In *Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1841–1857.
- Wolfram, W. and Eisenhauer, K. (2019). Implicit sociolinguistic bias and social justice. In *The Routledge Companion to the Work of John R. Rickford*, pages 269–280. Routledge.