



Universidade do Minho



UNIVERSIDADE
DE VIGO

*lingua*MATICA

Volume 18, Número 1 (2026)

ISSN: 1647-0818

lingua

Volume 18, Número 1 – 2026

LinguaMÁTICA

ISSN: 1647-0818

Editores Executivos

Marcos Garcia

Hugo Gonçalo Oliveira

Editores

Alberto Simões

José João Almeida

Xavier Gómez Guinovart

Conteúdo

Artigos de Investigação

A Transcrição Fonética como Ferramenta Pedagógica no Ensino de Português para Hispanofalantes: Um Estudo De Caso

Keila Coelho 3

Norma, ús i interferència: biaixos lingüístics en els models de llenguatge en català

Mireia Almena Rodríguez & Thomas Brochhagen 15

Um Comitê de Classificadores para Anotação Automática de Linguagem Tóxica em Português sob Escassez de Dados

F. A. R. Neto, R. T. Anchiêta, R. S. Moura, P. A. S. Neto & A. M. Santana 27

Revisão

A comissão científica da **LinguaMÁTICA** pode ser consultada na página Web da revista, em <https://linguamatica.com/index.php/linguamatica/about/editorialTeam>.

Para esta edição, colaboraram os seguintes investigadores:

- **Adelina Castelo**, Universidade Aberta
- **Aline Bezerra**, Universidade Federal de Alagoas
- **Antoni Oliver González**, Universitat Oberta de Catalunya
- **Alexandre Rademaker**, Fundação Getúlio Vargas
- **Eugénio Ribeiro**, Instituto Universitário de Lisboa
- **Francielle Alves Vargas**, Universidade de São Paulo
- **Gerardo Sierra**, Universidad Nacional Autónoma de México
- **Iria da Cunha**, Universidad Nacional de Educación a Distancia
- **Isabel Espinosa Zaragoza**, Universitat d’Alacant
- **Ivandré Paraboni**, Universidade de São Paulo
- **Jéssica Rodrigues da Silva**, University of Oxford
- **Jorge Baptista**, Universidade do Algarve
- **Lluís Padró**, Universitat Politècnica de Catalunya
- **Nora Graichen**, Universitat Pompeu Fabra
- **Rafael Anchiêta**, Instituto Federal do Piauí
- **Viviane Moreira**, Universidade Federal do Rio Grande do Sul

Artigos de Investigação

A Transcrição Fonética como Ferramenta Pedagógica no Ensino de Português para Hispanofalantes: Um Estudo De Caso

Phonetic Transcription as a Pedagogical Tool in Teaching Portuguese for Spanish Speakers: A Case Study

Keila Coelho  

Universidade de Sevilla

Resumo

Este estudo investiga o uso da transcrição fonética como ferramenta pedagógica no ensino de Português, com foco em falantes nativos de espanhol. A pesquisa parte do pressuposto de que dificuldades na pronúncia de determinados fonemas inexistentes na língua materna comprometem a aquisição da oralidade em português. O objetivo é avaliar em que medida a incorporação da transcrição fonética pode auxiliar na superação desses desafios, promovendo maior precisão na pronúncia e desenvolvimento da memória auditiva. A metodologia adotada foi qualitativa, de caráter descritivo e exploratório, com aplicação na Universidade de Sevilla (Espanha), entre 2022 e 2023. Participaram 46 estudantes espanhóis, distribuídos entre os níveis A2 e B1 do Quadro Europeu Comum de Referência para Línguas. Foram realizadas atividades com poemas, trava-línguas e textos literários, priorizando sons não existentes na língua materna dos alunos. A transcrição fonética foi introduzida progressivamente, com foco em palavras-chave. Os resultados indicam que os alunos que utilizaram essa estratégia demonstraram avanços significativos na pronúncia e na consciência fonológica, embora obstáculos iniciais tenham sido observados quanto à familiarização com o Alfabeto Fonético Internacional. Conclui-se que a transcrição fonética, quando usada de forma direcionada e progressiva, é uma ferramenta eficaz para o ensino de português, especialmente no desenvolvimento da competência oral. Sugere-se sua ampliação em contextos didáticos semelhantes, respeitando a complexidade e o tempo de assimilação dos estudantes.

Palavras chave

ensino de português; pronúncia; transcrição fonética

Abstract

This study investigates the use of phonetic transcription as a pedagogical tool in the teaching of Portuguese, with a focus on native Spanish speakers. The research is based on the assumption that difficulties in the pronunciation of certain phonemes that do not exist in the mother tongue jeopardize the acquisition of orality in Portuguese. The aim is to as-

sess the extent to which the incorporation of phonetic transcription can help overcome these challenges, promoting greater accuracy in pronunciation and the development of auditory memory. The methodology adopted was qualitative, descriptive and exploratory, and was carried out at the University of Seville (Spain) between 2022 and 2023. Forty-six Spanish students took part, distributed between levels A2 and B1 of the Common European Framework of Reference for Languages. Activities were carried out with poems, tongue twisters and literary texts, prioritizing sounds that do not exist in the students' mother tongue. Phonetic transcription was introduced progressively, focusing on key words. The results indicate that the students who used this strategy showed significant progress in pronunciation and phonological awareness, although initial obstacles were observed in terms of familiarization with the International Phonetic Alphabet. We conclude that phonetic transcription, when used in a targeted and progressive way, is an effective tool for teaching Portuguese, especially in developing oral competence. It is suggested that it be expanded in similar teaching contexts, respecting the complexity and assimilation time of the students.

Keywords

Portuguese language teaching; pronunciation; phonetic transcription

1. Introdução

A aquisição da oralidade em uma língua não materna exige muito mais do que o domínio das palavras e normas da gramática. Num livro recente, o linguista John Levin afirma que a evolução do ensino de línguas estrangeiras, ao longo dos últimos anos, tem vindo a consolidar uma série de ideias erradas sobre a aprendizagem da pronúncia. Para este autor, ideias como a de que a pronúncia é adquirida por si mesma, que os professores não podem ensinar pronúncia ou que os adultos não podem adquirir a pronúncia dos falantes nativos estão a ocupar cada vez mais espaço. Contra essas ideias, Levin defende que o



DOI: 10.21814/lm.18.1.494

This work is Licensed under a

Creative Commons Attribution 4.0 License

ensino da pronúncia constitui um elemento essencial no ensino de línguas estrangeiras e que, por meio de estratégias adequadas, pode funcionar, de forma variável, em alunos de qualquer idade (Levin, 2022).

A pronúncia apropriada dos fonemas constitui um fator fundamental para comunicação efetiva, como sublinham Neto & de Oliveira (2023). A correção fonética é necessária quando, na produção oral, são detectados erros que precisam ser corrigidos. Para esse fim, existe uma série de estratégias — descritas a seguir — que podem contribuir para a redução de problemas comumente agrupados sob a denominação de “sotaque estrangeiro”. Para falantes nativos de espanhol, o aprendizado do idioma português como língua não materna evidencia dificuldades fonológicas específicas, considerando que a língua portuguesa tem fones que não existem no espanhol, como as vogas nasais e certas consoantes fricativas, de acordo com aqueles autores.

Várias pesquisas indicam que a instrução explícita das normas fonológicas não é o bastante para assegurar uma pronúncia efetiva por parte dos alunos, tal e como indicam Telles & Lung (2024). Numa etapa inicial, a fonética contrastiva, quando utilizada de maneira adequada, pode constituir uma ferramenta valiosa para a previsão de erros. No entanto, a partir do momento em que o professor entra em contato direto com seus alunos, torna-se imprescindível a análise da interlíngua por eles desenvolvida. Tradicionalmente, os métodos de ensino de línguas estrangeiras, incluindo o português, têm se baseado na escuta e imitação de gravações ou do discurso do próprio docente, partindo do pressuposto de que, ao ouvir repetidamente determinados sons, o aprendiz seria capaz não apenas de reproduzi-los corretamente, mas também de integrá-los ao novo sistema fonético-fonológico da língua-alvo que está gradualmente construindo. Contudo, tal abordagem encontra um obstáculo conhecido como surdez fonológica. Este conceito, proposto pelo linguista Evgueni Polivanov em 1931, defende que “quando ouvimos uma palavra estrangeira desconhecida (ou, de uma maneira geral, um fragmento de língua estrangeira que, devido a seu volume, pode ser captado pela percepção auditiva), tratamos de reencontrar nela um complexo de representações fonológicas nossas, de decompô-la em fonemas peculiares à nossa língua materna e em conformidade até com nossas leis de agrupamento dos fonemas” (Polivanov, 1978). Assim o falante tende a não perceber e, conseqüentemente, a não realizar os fones que não existem em sua língua materna. Diante disso, o ensino da pronúncia

passa a depender de um processo de reeducação da percepção auditiva, com o objetivo de levar o aluno a assimilar de forma adequada as categorias fonéticas da língua estrangeira e, assim, alcançar uma pronúncia o mais próximo possível da nativa, de acordo com Listerri (2003). Mas convém sublinhar, tal e como indica Castelo (2018), que, atualmente, ao contrário do que se procurava no sistema tradicional de ensino da pronúncia, o objetivo deste ensino é alcançar um bom nível de qualidade na expressão oral, mas não necessariamente a imitação da pronúncia nativa. Assim, tal como consideram Derwing & Murray (2009), o sotaque ao utilizar uma língua estrangeira pode servir de base para fazer julgamentos de caráter social, mas não tem por que comprometer a comunicação. Somente nos casos em que a inteligibilidade é dificultada é que a intervenção de um professor é absolutamente necessária.

Neste contexto, e ao longo de um período relativamente recente, surgiram novas propostas que procuram superar as limitações acima mencionadas através de métodos pedagógicos multidisciplinares. Entre elas, destaca-se o método MMA (*Multisensory, Multicognitive Approach*), proposto por Odisho (2007), que se baseia numa abordagem multissensorial e multicognitiva do ensino da pronúncia, envolvendo recursos auditivos, visuais e motores. Assim, para este autor, os alunos devem ser incentivados não apenas a ouvir os sons da língua estrangeira, mas também a memorizá-los a curto prazo e a contrastá-los com os sons que fazem parte do seu inventário psicolinguístico, ou seja, os da sua língua nativa.

Neste sentido, o uso de recursos visuais e sonoros que estimulem a memória auditiva e a consciência articulatória dos discentes vem se mostrando uma forma eficiente de superar esses desafios. Nesse cenário, a transcrição fonética, acima de tudo através do Alfabeto Fonético Internacional (AFI), permite estabelecer elementos visuais que atuam como alertas durante a leitura em voz alta de textos escritos, especialmente em relação aos fonemas que exigem atenção especial por parte do aluno. Uma função desses alertas é atenuar a surdez fonológica. Outra função é a de reduzir o impacto da interlíngua, entendida neste contexto como a interferência do sistema fonético da língua materna na aquisição de uma língua estrangeira, de acordo com Alexopolou (2010). Dessa forma, essa transcrição desafia o aluno a relacionar os símbolos fonéticos aos sons do português, podendo contribuir para uma pronúncia mais consciente e precisa.

O presente trabalho apresenta como problemática pesquisar de que maneira a transcrição fonética pode contribuir para o aprimoramento da pronúncia de discentes hispanofalantes no processo de aprendizagem do português como língua não materna. Parte-se da hipótese de que a inserção da transcrição fonética diretamente nos textos lidos em voz alta pelos estudantes favorece o desenvolvimento da competência oral, especialmente no que se refere à distinção e produção de fonemas desafiadores. O objetivo principal da pesquisa é analisar a eficácia da transcrição fonética como ferramenta pedagógica no ensino de Português como Língua Não Materna (PLNM). Os objetivos secundários são: identificar os principais fonemas que causam dificuldades aos hispanofalantes; aplicar atividades com uso da transcrição fonética em diferentes gêneros textuais; e compreender os efeitos dessa prática na pronúncia e na consciência fonológica dos estudantes, entendida como um conjunto de competências metalinguísticas que permitem ao aluno analisar, processar e compreender os elementos fonêmicos para posterior produção oral, separando as palavras em elementos mais pequenos, como as sílabas, como sublinha [Fernández \(2013\)](#).

O artigo se justifica pela necessidade de estudar metodologias que vão além da explicação teórica e do sistema de imitação e que atendam às particularidades fonológicas de aprendizes hispanofalantes. A transcrição fonética, quando utilizada de forma gradual e direcionada, pode suprir lacunas de ensino, permitindo ao aluno uma aprendizagem mais significativa e articulada com as demandas da oralidade. A pesquisa foi realizada na Universidade de Sevilha, na Espanha, durante o ano acadêmico 2022/2023. Adotou-se uma abordagem qualitativa, com caráter descritivo e exploratório, envolvendo 46 estudantes espanhóis, com idades entre 18 e 46 anos, distribuídos nos níveis A2 e B1 segundo o Quadro Europeu Comum de Referência para Línguas ([Conselho de Europa, 2001](#)). As atividades didáticas incluíram poemas, trava-línguas e textos literários, com foco em fonemas inexistentes no espanhol, e a transcrição fonética foi aplicada em etapas, sempre com suporte do professor e prática oral orientada.

2. O sistema vocálico e consonantal no Português e no Espanhol: semelhanças, diferenças e implicações didáticas

A maior parte dos teóricos que estudam as estratégias de ensino de línguas estrangeiras coincidem em que o estudo do sistema fonológico dos idiomas é fundamental para o entendimento dos obstáculos encontrados por alunos de uma outra língua, segundo [Stein \(2011\)](#). No caso dos hispanofalantes que aprendem português, o sistema consonantal evidencia tanto pontos de aproximação como de expressiva divergência entre os idiomas, o que traz impactos diretos na aprendizagem da pronúncia certa. Mesmo que a língua portuguesa e a espanhola compartilhem uma quantidade significativa de fonemas consonantais, é fundamental ressaltar que a forma como esses sons são articulados, distribuídos e percebidos muda expressivamente entre os dois sistemas linguísticos, como salienta [Cristófar-Silva \(2019\)](#). No [Tabela 1](#) que apresenta um quadro de consoantes, é notória a existência de fonemas consonantais em comum entre o português e o espanhol.

Dentre os fonemas comuns a ambos os idiomas, estão as oclusivas /p/, /t/, /k/, /b/, /d/, /g/, as nasais /m/ e /n/ e a lateral /l/, que normalmente não evidenciam grandes obstáculos para os alunos hispanofalantes. Porém, esses fo-

Modo de articulação	Ponto de articulação	Português	Espanhol
Oclusivas	Bilabial	/p, b/	/p, b/
	Alveolar	/t, d/	/t, d/
Fricativas	Velar	/k, g/	/k, g/
	Labiodental	/f, v/	/f, v/
Fricativas	Labiodental	/f, v/	/f/
	Alveolar	/s, z/	/s/
	Pós-alveolar	/ʃ, ʒ/	/j/ (algumas regiões)
Nasais	Velar	/x (alguns dialetos)	/x/
	Bilabial	/m/	/m/
Nasais	Alveolar	/n/	/n/
	Palatal	/ɲ/	/j/
Laterais	Alveolar	/l/	/l/
	Palatal	/j/	/j/ (zonas não yeístas)
Vibrantes	Alveolar (simples)	/r/	/r/
Alveolar (simples)	Alveolar (múltipla)	/r/	/r/
Palatal	Palatal	/j/	/j/
Aproximantes	/j/	/j/	/j/ (zonas não yeístas)
Palatal	Labio-velar	/w/	/w/

Tabela 1: Contrastivo simplificado dos sons consonantais do português e do espanhol. Fonte: Adaptado para os sons do português de [Freitas et al. \(2012\)](#), e para os do espanhol de [Gil \(2007\)](#).

nemas podem evidenciar variações fonéticas notáveis, principalmente no que se refere à tonicidade, contexto silábico e entonação (Freitas et al., 2012). Ademais, existem discrepâncias consideráveis que atingem diretamente a fluência e inteligibilidade dos falantes de espanhol em português.

Uma das disparidades mais expressivas está nas fricativas alveopalatais surdas e sonoras, /ʃ/ (como em *chuva*) e /ʒ/ (como em *joia*), existentes na língua portuguesa, porém inexistentes na língua espanhola padrão. Em diversos casos, alunos hispanofalantes costumam trocar /ʃ/ por /s/ ou /tʃ/, e /ʒ/ por /j/, como consequência de uma transferência negativa do idioma materno. Essa condição, chamada de interferência fonológica, pode prejudicar a transparência da fala e causar problemas semânticos, principalmente em pares mínimos, tal e como sublinha Lipski (2008). Outra questão importante diz respeito à realização do fonema /r/. No português brasileiro, esse fonema pode ser realizado, em posição intervocálica ou inicial, como uma fricativa uvular ou glotal, em palavras como *terra* e *rei*. No português europeu, o fonema /r/ é também frequentemente realizado como fricativa uvular /R/. Já no espanhol, não ocorre a fricativa uvular /R/, sendo o fonema /r/ realizado como vibrante alveolar. Por sua vez, o segmento /r/, um "tap" alveolar, está presente no português brasileiro, no português europeu e no espanhol. A complexidade de produção desse fonema faz com que diversos alunos produzam o som de maneira restrita ou inadequada, o que pode prejudicar tanto a inteligibilidade como a naturalidade da pronúncia, de acordo com Gomes (2021).

O fonema lateral palatal /ʎ/, representado na língua portuguesa pelo dígrafo "lh", como em *ilha*, constitui outro problema importante. Enquanto alguns dialetos espanhóis têm um som parecido, diversos outros adotaram o *yeísmo*, unindo os fonemas /ʎ/ e /j/. Consequentemente, alunos provenientes de nações como Colômbia, Argentina, México e Chile costumam trocar o som de "lh" por /j/, atrapalhando a identificação e a reprodução apropriada desse fonema, tal e como foi verificado por Alves et al. (2021) num estudo. Essas disparidades fonológicas, mesmo que diversas vezes sutis para os falantes nativos, possuem repercussão direta na compreensão e no desenvolvimento da fala dos alunos.

Depois do estudo dos sistemas consonantais dos dois idiomas, também é importante entender as disparidades entre os sistemas vocálicos da língua portuguesa e espanhola, considerando que essas disparidades atingem consideravelmente a inteligibilidade da fala dos alunos hispanofalantes.

	Anterior	Central	Posterior
Alta	/i/		/u/
Média	/e/		/o/
Baixa		/a/	

Tabela 2: Sons vocálicos de espanhol. Fonte: Adaptado de Rodríguez (2004).

Vogais orais			
Altura \ Posição	Anterior	Central	Posterior
Alta	/i/	ĩ/	/u/
Média-alta	/e/	—	/o/
Média-baixa	/æ/	—	/σ/
Baixa	—	/a/	

Vogais nasais			
Altura \ Posição	Anterior	Central	Posterior
Alta	ĩ/	—	/ú/
Média	/è/	—	/ô/
Baixa	—	a/	

Tabela 3: Sons vocálicos do português. Fonte: Adaptado de João et al. (2024).

A Tabela 2 mostra o sistema vocálico da língua espanhola, caracterizado por um inventário limitado de somente cinco vogais orais: /i, e, a, o, u/. Esse sistema fonológico relativamente simples, fundamentado somente na diferença entre altura e anterioridade/posterioridade, colabora com uma maior estabilidade fonética na língua espanhola e restringe a ocorrência de variações vocálicas em contextos fonológicos diferentes.

Já a Tabela 3 apresenta o sistema vocálico da língua portuguesa, que se difere pela maior dificuldade e número de vogais, de acordo com Velloso (2012). Enquanto a língua espanhola possui somente cinco vogais orais, a variedade brasileira do português tem sete vogais orais plenas (/i, e, ε, a, ɔ, o, u/) e cinco vogais nasais fonêmicas (/ĩ, ê, ɐ̃, õ, ã/), além de variações regionais e processos de redução vocálica em sílabas átonas, tal e como indica Rodríguez (2004). Essa dificuldade é um dos elementos que potencialmente podem dificultar a pronúncia de hispanofalantes que estão aprendendo português porque diversas dessas vogais não têm correspondentes diretos no espanhol.

Um outro desafio importante encontrados pelos falantes de espanhol ao aprenderem a língua portuguesa se refere à diferenciação entre vogais médias abertas e fechadas. O português diferen-

cia foneticamente pares como /e/ e /ɛ/ (existentes em palavras como *selo* vs. *sela*), e /o/ e /ɔ/ (*avô* vs. *avó*), o que não acontece no espanhol, onde essa distinção não é fonológica. Por conseguinte, alunos hispanofalantes constantemente pronunciam as duas variantes como um único som médio (/e/ = /ɛ/, /o/ = /ɔ/), o que pode prejudicar o entendimento e causar confusões semânticas, tal e como considera Neto (2014).

3. A Transcrição Fonética como ferramenta pedagógica no ensino de Português como Língua Não Materna

A transcrição fonética, fundamentada no AFI, é bastante conhecida e valorizada como um recurso didático importante para o ensino de idiomas estrangeiros, principalmente no que se refere à pronúncia. Celce-Murcia et al. (2010) destacam que, ao representar de forma gráfica os sons da fala de forma precisa e padronizada, ela possibilita que os alunos enxerguem e articulem adequadamente os fones que não pertencem a seu idioma materno. Na abordagem do ensino do PLNM, esse recurso vem demonstrando potencial para aumentar a consciência fonológica dos alunos, além de colaborar com a criação de uma memória auditiva bastante aprimorada.

Alguns teóricos como Neto & de Oliveira (2023) discutem que a transcrição fonética promove apoio visual para a articulação adequada dos sons, acima de tudo quando relacionada a atividades de leitura em voz alta e escuta atenta. Estas atividades possibilitam que o aluno relacione grafemas e fonemas de maneira mais efetiva, o que acarreta mecanização da pronúncia exata, principalmente em idiomas mais complexos fonologicamente, como a língua portuguesa. Ademais, pesquisas como a de Corrêa et al. (2010) mostram que a utilização sistemática da transcrição fonética pode diminuir influências do idioma materno e colaborar com a independência dos discentes na correção de sua própria pronúncia.

A efetividade desse enfoque está relacionada à maneira com o docente faz o papel de mediador do processo de aprendizagem fonológica. De acordo com Cristófar-Silva (2019), a mera apresentação dos símbolos do AFI não é o bastante: é preciso uma atividade didática cautelosa, que leve em consideração a cadência, entonação, as características suprasegmentais e o contexto de utilização do idioma. Dessa forma, a transcrição fonética precisa ser inserida de maneira gradativa, privilegiando sons que apresentem maior nível de dificuldade e usando exemplos verdadeiros

retirados de textos consideráveis, como poesias, trava-línguas e texto originais da cultura-alvo.

3.1. Dificuldades fonológicas específicas de hispanofalantes no aprendizado do português

Apesar das semelhanças entre o português e o espanhol — línguas românicas com estruturas fonológicas próximas —, existem diferenças significativas no inventário de sons que afetam diretamente a aquisição da pronúncia por falantes nativos de espanhol. Essas divergências constituem um dos principais obstáculos no desenvolvimento da competência oral em português. Uma das maiores dificuldades está associada à existência de vogais nasais no português, como [ẽ], [ē], [ō], que não existem no espanhol. Essa falta torna complexo para os alunos o entendimento auditivo e o desenvolvimento apropriado dessas vogais, constantemente levando à sua oralidade inadequada.

Outra questão de diferença fonológica importante são as fricativas palatais /ʃ/ e /ʒ/, de acordo com Silveira & Souza (2011), frequentes em português, porém ausentes no espanhol. Falantes de espanhol costumam trocar /ʃ/ por /s/ ou /tʃ/, e /ʒ/ por /j/ ou /ʒ/, o que atrapalha a inteligibilidade e pode causar distorções semânticas, como sublinha Lipski (2008). A produção do “r” intervocálico e inicial também evidencia um desafio. No português brasileiro, o /r/ em posição inicial ou intervocálica pode ser realizado como fricativa velar ou uvular surda [x], bastante diferente do “r” vibrante simples frequente no espanhol. Essa diferença promove constantemente a hipercorreção ou omissão dos sons por parte dos aprendizes, tal e como indica Cristófar-Silva (2019).

Ademais, a língua portuguesa possui uma maior variação de timbres vocálicos, somando até 7 fonemas vocálicos, que se concretizam em 14 fones vocálicos, enquanto a língua espanhola atua com somente 5 vogais orais plenas. Estes contrastes fonológicos manifestam-se em pares mínimos, por exemplo contrastando /e/ e /ɛ/, ou /o/ e /ɔ/, acarretando pronúncias que, mesmo compreensíveis, são caracterizadas por intenso sotaque e pouca precisão, de acordo com Perozzo & Kupske (2024). Sons como o /ɰ/ e o /j/ evidenciam variações expressivas nos dialetos hispânicos. Em alguns países latino-americanos, o /ɰ/ quase não existe, sendo trocado por /j/, o que causa outros obstáculos de adequação fonológica na língua portuguesa, como sublinham Adamczyk et al. (2021).

4. Metodologia

A presente pesquisa foi conduzida na Universidade de Sevilla (Espanha) durante o ano acadêmico 2022/23. Trata-se de um estudo de natureza qualitativa, com caráter descritivo e exploratório, cujo objetivo principal foi analisar a eficácia da transcrição fonética como ferramenta pedagógica no ensino de PLNM, com foco específico nos desafios enfrentados por falantes nativos de espanhol. A investigação foi realizada em um contexto universitário real de aprendizagem, buscando compreender como os estudantes lidam com os fonemas ausentes em sua língua materna e de que maneira a transcrição fonética pode auxiliar nesse processo. A amostra foi composta por 46 estudantes, falantes da variedade do espanhol da Espanha e com idades variando entre 18 e 46 anos, matriculados em cursos de português como língua adicional. Os participantes estavam distribuídos entre os níveis A2 e B1 do Quadro Europeu Comum de Referência para Línguas, de modo a representar alguma diversidade tanto etária quanto de proficiência linguística. A seleção dos participantes buscou contemplar diferentes perfis de aprendizes, o que possibilitou uma análise mais abrangente dos efeitos da metodologia aplicada. O período da intervenção didática foi de quatro semanas, distribuídas em períodos de uma semana para cada uma das atividades didáticas descritas logo a seguir.

4.1. Atividades Didáticas Aplicadas

A intervenção pedagógica foi dividida em etapas progressivas, visando a introdução gradual da transcrição fonética e o treinamento auditivo e articulatório dos aprendizes. A primeira etapa consistiu em um estudo comparativo entre os fonemas do espanhol e do português, com foco especial nos sons que não possuem equivalência direta entre as línguas. O modelo de pronúncia utilizado foi o da variedade brasileira do português. Esse momento permitiu que os alunos tomassem consciência das diferenças fonológicas fundamentais e estabelecessem relações entre os dois sistemas.

Na segunda fase, foi aplicada uma atividade com o poema “Bolhas”, de Cecília Meireles. Os estudantes ouviram a leitura do poema, atentando-se à melodia e à sonoridade das palavras. Em seguida, participaram de uma discussão coletiva, identificando sons que lhes pareciam estranhos ou distintos do espanhol. O professor destacou vogais nasais (como em *mundo*), vogais orais abertas e fechadas (como em *bolhas*), e consoantes desafiadoras como o “r” inicial de *rodar*. As palavras selecionadas foram transcritas foneticamente

com auxílio do professor, utilizando os símbolos do AFI:

- *bolhas* [ˈboʎɐʃ].
- *mundo* [ˈmũdu].
- *rodar* [ʁoˈdaɾ].

Essa prática foi acompanhada de feedback individualizado, o que possibilitou a correção dos erros de pronúncia de forma direcionada.

A atividade seguinte trabalhou o poema “Canção do Exílio”, de Gonçalves Dias, escolhido pela presença marcante de vogais nasais, consideradas um dos maiores desafios para os aprendizes. Após contextualização histórica do poema, foi realizada uma leitura expressiva, seguida pela identificação de palavras com sons complexos como *minha*, *tem*, *cantam* e *lá*. Também foram comparadas as vogais nasais com suas possíveis aproximações em espanhol, e a articulação do “r” gutural foi enfatizada. Palavras-chave foram transcritas com a ajuda dos estudantes:

- *minha* [ˈmijɐ].
- *tem* [tẽj̃].
- *rodar* [ʁoˈdaɾ].

Além disso, foram destacados os seguintes pontos fonológicos críticos para os hispanofalantes:

1. Nasalização: [ẽ], [ɐ̃], [õ] – em palavras como *tem*, *canta*, *não*;
2. Som do “r”: [r] – como em *terra*, *morra*;
3. Ditongos nasais: como em *mãe* [mẽ̃j̃];
4. Vogais abertas e fechadas: como em *céu* [sɛw] e *eu* [ɛw].

Esses sons foram comparados com palavras semelhantes ou iguais em espanhol (ex: *céu/cielo*, *flores/flores*) para facilitar a compreensão articulatória e semântica.

Outra atividade consistiu na leitura e prática de trava-línguas tradicionais, adaptados aos grupos, com foco nos sons mais desafiadores para falantes de espanhol. Os alunos formaram pares e leram os trava-línguas alternadamente, primeiro em ritmo lento, com atenção à pronúncia correta, e depois em ritmo natural, seguindo o modelo de pronúncia da variedade brasileira do português. Os textos selecionados incluíram:

1. *O rato roeu a roupa do rei de Roma* • [ˈɾatu], [ʁuˈɛw], [ˈʁowɐ], [ˈɾɛj], [ˈʁomɐ] → Desafio: som [ʁ] no início das palavras.

2. *Três pratos de trigo para três tigres tristes* • [treˈʃ], [pratuʃ], [trigu], [tigrɪʃ], [triʃtiʃ] → Desafio: encontros [tr], [pr].
3. *A aranha arranha a jarra, a jarra arranha a aranha* • [ɐˈʁɐɾɐ], [ʒaʁɐ] → Desafio: vibrante [ʁ], nasal [ɹ] e fricativa [ʒ].

Essas atividades foram acompanhadas de repetições em coro e práticas individuais, seguidas por momentos de partilha das dificuldades e descobertas por parte dos estudantes. A articulação de sons problemáticos foi reforçada com apoio visual (transcrição fonética) e auditivo, visando consolidar a memória sonora e melhorar a fluência oral em português.

A seguinte atividade desenvolvida envolveu o poema “A mulher mais bonita do mundo” (2002), de José Luís Peixoto, pertencente ao universo do português europeu. Essa escolha teve como objetivo familiarizar os estudantes com as variações fonéticas entre o português brasileiro e o europeu, destacando como a pronúncia e a articulação dos fones pode variar entre as variedades da língua portuguesa. O poema, lido pela professora, foi o seguinte:

*Estás tão bonita hoje. quando digo que
nasceram flores novas na terra do jar-
dim, quero dizer que estás bonita. entro
na casa, entro no quarto, abro o armá-
rio, abro uma gaveta, abro uma caixa
onde está o teu fio de ouro. [...]*

Na primeira parte da atividade, os alunos analisaram desde o ponto de vista teórico os sons com o apoio da professora. Palavras-chave foram destacadas no texto, e os fonemas mais desafiadores foram sublinhados em suas cópias impressas. A professora, após a leitura do poema, enfatizou os sons a serem trabalhados. Os alunos, então, identificaram os fonemas mais difíceis, como:

- “*estás*” – foco no som final [ʃ].
- “*tão*” – atenção ao ditongo nasalizado [ɐ̃w̃].
- “*hoje*” – prática do som fricativo palatal sonoro [ʒ].

Na segunda parte, os alunos participaram de uma repetição guiada. Eles escutaram a leitura feita pela professora ou gravações e repetiram as palavras destacadas, praticando isoladamente os sons mais complexos:

- Vogais nasais: como em *tão* [tɐ̃w̃], *mundo* [ˈmũdu].
- Fricativas palatais: como em *hoje* [ˈoʒɪ], *jardim* [ʒarˈdĩ].

- Vogais reduzidas: como em *bonita* [boˈnɪtɐ], *beleza* [beˈlezɐ].

Por fim, os alunos compararam sua própria pronúncia com a transcrição fonética fornecida e identificaram os pontos que necessitavam de maior prática. Essa etapa de autocorreção foi essencial para o desenvolvimento da consciência fonológica e para o fortalecimento da autonomia do aprendiz no processo de correção da pronúncia. Para garantir um aprendizado significativo e evitar a sobrecarga cognitiva, o material fonético foi dividido em partes menores, com foco na familiarização dos alunos com os sons-alvo, sem exigir memorização imediata. Essa abordagem favoreceu a assimilação natural dos fonemas e preparou os estudantes para um domínio mais apurado da articulação oral.

Como exercício final do conjunto de atividades, foi aplicado um trecho retirado de uma reportagem do *Jornal do Brasil*, intitulada “Ciência precisa dialogar com ecologia indígena, dizem pesquisadores” fragmento selecionado oferece uma oportunidade relevante para a prática da pronúncia de palavras de uso formal e científico, com foco em sons complexos para falantes hispanofalantes, como vogais nasais, encontros consonantais e variações de entonação:

*É preciso a superação [superaˈsɐ̃w̃]
dessa compreensão ciência [ˈsjɛ̃sjɐ]-
natureza, cultura-natureza. Enquanto
os seres humanos entendem que os ou-
tros seres são matéria, que não pensam
como nós, humanos, isso [o desequilíbrio
[dezekiˈlibrju] ambiental] vai continuar
existindo. Os indígenas [iˈɕiʒenɐs] têm
outra visão”, aponta Justino Sarmiento.
(Jornal do Brasil, 2024)*

Essa atividade teve como objetivo reforçar a leitura oral com o apoio da transcrição fonética e estimular a comparação fonológica com o espanhol, por meio de um gênero textual distinto (jornalístico e argumentativo), ampliando o repertório linguístico dos alunos. As etapas propostas foram as seguintes:

1. Leitura em voz alta do texto, com atenção à pronúncia das palavras destacadas e à transcrição fonética fornecida pelo professor.
2. Comparação com o espanhol: os alunos foram convidados a refletir sobre como essas palavras seriam pronunciadas em sua língua materna e a identificar as diferenças mais significativas, especialmente:

- (a) Nos sons nasais: *superação* [supɐɾa'sɐ̃w̃], *ciência* ['sjɛ̃sja], *indígenas* [ĩ'ɕizɐnɐ̃s].
 - (b) Nas consoantes fricativas sonoras: *desequilíbrio* [dezɛki'libɾju].
 - (c) Na entonação de enumerações e pausas reflexivas, comum no estilo argumentativo.
3. Os estudantes também praticaram a autocorreção fonética, comparando suas produções com a pronúncia-modelo realizada oralmente pela professora e refletindo sobre os ajustes necessários. A transcrição fonética foi usada como suporte visual e articulatório, funcionando como um mapa para guiar a emissão dos sons.

Essa atividade possibilitou o reforço de aprendizagens anteriores, em especial no que se refere à nasalização, à identificação de vogais médias abertas, à produção de fricativas sonoras palatais (como [ʒ]), e à segmentação correta de sílabas complexas. O uso de um texto não literário também proporcionou diversidade de registros linguísticos, contribuindo para a competência comunicativa dos alunos em contextos mais formais.

4.2. Procedimentos da recolha e análise de dados

Os dados da presente pesquisa foram coletados através da observação direta e anotações do estudo de caso, produzidos ao longo das aulas nas quais os exercícios fonéticos foram desenvolvidos. A investigadora, que também exerceu a função de docente ao longo das aulas, registrou condutas linguísticas, desafios encontrados com frequência, reações e autoanálises dos discentes e avanço na pronúncia oral. Além da análise, foram observadas gravações de áudio das leituras em voz alta feitas no começo e no fim de cada módulo de exercícios. Essas anotações possibilitam equiparar a performance dos discentes em relação à precisão articulatória antes e depois da inserção da transcrição fonética. Para a análise dos dados, foi usada a análise qualitativa de conteúdo, fundamentada em categorias antecipadamente estabelecidas a partir da finalidade do trabalho. As categorias foram:

1. Desafios fonológicos específicos;
2. Desempenho na produção dos sons-alvo;
3. Compreensão fonética dos estudantes;
4. Resultados da transcrição fonética na autonomia do aluno.

Com essas categorias, as informações foram estruturadas, codificadas e compreendidas de acordo com as referências bibliográficas da área, procurando encontrar padrões, progressos e limitações no processo de aprendizagem fonológica. A apresentação e discussão dos dados serão organizados em torno das quatro categorias supramencionadas. Para cada uma delas, serão mostrados exemplos retirados dos exercícios desenvolvidos, assim como excerto de falas dos aprendizes, que evidenciam seu progresso e compreensão acerca das próprias dificuldades fonológicas. Serão inseridas transcrições fonéticas comparadas (começo e fim) de palavras-chave trabalhadas no decorrer das aulas, permitindo uma leitura transparente dos resultados da prática com transcrição fonética. A discussão será fundamentada em teóricos do campo da fonética aplicada e do ensino de idiomas, articulando os dados empíricos com a literatura usada no estado da arte da pesquisa.

5. Resultados e Discussão

No que diz respeito aos resultados obtidos cabe destacar que a análise quantitativa dos dados revelou avanços consideráveis na produção dos fonemas críticos pelos estudantes desde o início até o final das intervenções com transcrição fonética. O gráfico abaixo sintetiza a comparação entre os desempenhos médios dos 46 participantes desde o começo até o final da prática didática orientada, em relação a fonemas que historicamente apresentam maior grau de interferência da língua materna espanhola: o lateral palatal /ʎ/, a fricativa uvular /ʁ/, a fricativa sonora palatal /ʒ/ e a vogal nasal /ɐ̃/.

Para fins de análise fonética dos resultados, adotou-se uma escala qualitativa de pontuação de 0 a 3, aplicada especificamente à avaliação da pronúncia dos fonemas críticos trabalhados ao longo das atividades. Esta escala permitiu avaliar a qualidade fonética observada na produção oral dos estudantes com base nos seguintes critérios: 0 – produção incorreta ou ininteligível; 1 – produção com erro sistemático ou forte interferência da língua materna, ainda que inteligível; 2 – produção parcialmente correta, com leve desvio de sotaque ou hesitação; 3 – produção correta, natural e foneticamente adequada. Essa abordagem de avaliação está alinhada a práticas consolidadas na fonética aplicada ao ensino de línguas estrangeiras, como descrevem [Celce-Murcia et al. \(2010\)](#), e foi escolhida por sua clareza e facilidade de aplicação no contexto de observação direta em sala de aula.

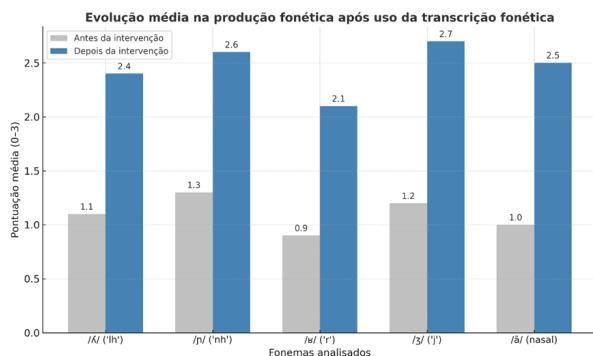


Figura 1: Evolução média na produção fonética após uso da transcrição fonética

Observa-se que todos os fonemas analisados apresentaram aumento significativo na média de produção fonética correta, segundo a avaliação realizada pela investigadora, sendo os mais expressivos:

1. O fonema /ɐ̃/, que evoluiu de 1,0 para 2,5, representando um ganho de 150% na produção correta do traço de nasalidade fonêmica.
2. O fonema /ɾ/, cuja média inicial de 0,9 passou para 2,1, indicando melhora de mais de 130% na articulação do “r” gutural, frequentemente inexistente no espanhol peninsular.
3. O fonema /ʒ/ também apresentou grande avanço (de 1,2 para 2,7), demonstrando que a prática com trava-línguas e textos literários foi eficaz na redução da confusão articulatória com sons próximos da L1, como /j/. O mesmo se verificou para /ɲ/ (de 1,3 para 2,6), cuja complexidade foi parcialmente superada com o apoio visual dos símbolos do AFI e da repetição orientada.

Entretanto, persistiram algumas dificuldades específicas, especialmente com o fonema /ɫ/ (como em “milho” e “palha”), cuja média final foi de 2,4 – ligeiramente inferior aos demais. Isso sugere que, apesar do progresso, a substituição sistemática por /j/ em alunos oriundos de regiões hispânicas com “yeísmo” ainda constitui um obstáculo fonético considerável.

Além disso, os dados indicam que os fonemas palatais exigem mais tempo de assimilação, e sua estabilização depende fortemente da repetição contextualizada e da autocorreção contínua. De modo geral, os resultados confirmam que a intervenção didática, como um todo, com transcrição fonética pode promover um avanço objetivo e mensurável na produção oral dos aprendizes de PLN, especialmente quando integrada a práticas significativas de leitura em voz alta, escuta atenta e mediação docente.

A possibilidade de perceber os sons através dos símbolos do AFI possibilitou que os alunos entendessem de maneira mais transparente as distinções entre os fonemas da língua portuguesa e da espanhola, proporcionando um enfoque mais reflexivo e ativo no processo de aprendizagem da oralidade. Portanto, os resultados encontrados no decorrer dos exercícios apontam que o uso da transcrição fonética como recurso pedagógico teve consequências positivas no desenvolvimento da habilidade oral dos alunos hispanofalantes aprendizes de língua portuguesa como segundo idioma. A introdução, no texto escrito, da transcrição de determinados fonemas segundo o AFI possibilitou criar no aluno uma sensação de alerta, que contribuiu para reduzir, em diferentes graus, o recurso à interlíngua.

Mas estes resultados não podem ser interpretados de um modo absoluto. Trata-se de um estudo de caráter exploratório, portanto, a falta de um grupo de controlo em que houvesse também uma intervenção didática, mas sem transcrição fonética, impede contrastar de modo individualizado o efeito de esta ferramenta sobre a aprendizagem dos alunos, o que constitui a principal limitação deste estudo.

Todavia, também foram encontradas outras limitações. A maior delas diz respeito à oralidade dos símbolos fonéticos do AFI, que primeiramente provocou dúvidas e resistência por parte de alguns alunos. Diversos deles revelaram o sentimento de intimidação frente ao novo sistema de escrita, diferente do alfabeto tradicional. Esse desafio, porém, foi aos poucos superado devido à evolução pedagógica adotada, com a inserção dos símbolos de maneira gradativa, sempre fundamentada em termos contextualizados e com auxílio de prática auditiva e articulatória guiada. A superação desses desafios mostra que o sucesso da transcrição fonética está bastante associado à mediação do professor e à maneira como os conteúdos fonéticos são inseridos no dia a dia das aulas. A familiarização com os símbolos não tem de acontecer isoladamente, mas articulada com textos expressivos, leituras em voz alta e comparação fonológica com o idioma materno do aprendiz. Quando conectada com essas atividades, a transcrição fonética evidenciou-se um recurso técnico e didático de elevado valor formativo no aprimoramento da pronúncia dos alunos.

Finalmente, cabe referir o período da intervenção docente, limitado a um mês, como um fator que impede desenvolver uma metodologia dirigida a conseguir resultados mais detalhados e objetivos.

6. Conclusão

O estudo realizado mostrou que, apesar das limitações descritas, a utilização sistemática da transcrição fonética, quando relacionada a práticas contextualizadas e conduzidas por mediação do professor de forma qualificada, pode colaborar de modo relevante com o aperfeiçoamento da pronúncia, do entendimento fonológico e da autonomia dos alunos no processo de correção oral. A maior contribuição desta pesquisa para a área científica da Educação está na demonstração de que a transcrição fonética pode ser mais do que uma ferramenta técnica restrita à formação linguística. Quando apropriadamente adequada ao cenário didático, ela se evidencia um recurso de mediação entre a oralidade e letramento fonológico, proporcionando uma aprendizagem mais articulada e consciente das estruturas sonoras da língua a ser aprendida.

A pesquisa salienta a necessidade de se inserir atividades fonéticas sistematizadas no currículo de PLNM, acima de tudo para alunos cujos idiomas maternos possuem diferenças fonológicas importantes relacionadas à língua portuguesa. Do ponto de vista pedagógico, os resultados desta pesquisa evidenciam a relevância de metodologias que integrem o uso do AFI com práticas de escuta, leitura em voz alta, produção oral e análise comparativa entre línguas, ressaltando que a introdução progressiva dos símbolos fonéticos, aliada ao trabalho com textos literários e jornalísticos, contribuiu para reduzir resistências iniciais e facilitar a assimilação dos sons-alvo pelos estudantes.

Contudo, a pesquisa apresenta limitações, como a aplicação em um único contexto institucional e a restrição da amostra aos níveis A2 e B1, o que limita a generalização dos resultados, além do tempo reduzido para observação longitudinal imposto pelo calendário acadêmico. Para estudos futuros, recomenda-se a ampliação para contextos multilíngues, a inclusão de aprendizes com outras línguas maternas, a investigação da eficácia da transcrição fonética em níveis mais avançados de proficiência e sua aplicação com o apoio de tecnologias digitais, como softwares de reconhecimento de fala, além da formação docente voltada ao uso crítico e criativo dessa ferramenta.

Referências

- Adamczyk, Magdalena, Marta Pitat, Ana Garrido & Martin Testa. 2021. *La enseñanza del español como lengua extranjera en el siglo XXI: desafíos y estrategias*. Instituto de Estudios Ibéricos e Iberoamericanos de la Universidad de Varsovia
- Alexopolou, Angelica. 2010. La función de la interlengua en el aprendizaje de lenguas extranjeras. *Revista Nebrija de Lingüística Aplicada* 9. 1–9. [↗](#)
- Alves, Thamy, Lucinda Teixeira & Maria Silva. 2021. Cartografia linguística: uma análise fonético-fonológico do fenômeno linguístico – transformação do /lh/ em /i/ na cidade de Belém –PA. *Caderno de Letras* 40. 471–486. [doi](#) 10.15210/cdl.v0i40.17003
- Castelo, Adelina. 2018. Ensino da componente fonético-fonológica: uma síntese e um exemplo de português para estrangeiros. *Revista de Estudos Linguísticos da Universidade do Porto* 12. 41–72. [↗](#)
- Celce-Murcia, Marianne, Donna Brindon & Janet Goodwin. 2010. *Teaching pronunciation: A course book and reference guide*. Cambridge University Press
- Conselho de Europa. 2001. *Quadro europeu comum de referência para as línguas: Aprendizagem, ensino, avaliação*. ASA Edições
- Corrêa, Leda, Manuela de Jesus & Daniele Dos Anjos. 2010. Pronúncia do português brasileiro e sua importância na elaboração de dicionário de português para estrangeiros. *Interdisciplinar* 12. 49–61. [↗](#)
- Cristófaros-Silva, Thais. 2019. *Fonética e fonologia do português: roteiro de estudos e guia de exercícios*. Contexto
- Derwing, Tracy & Munro Murray. 2009. Putting accent in its place: Rethinking obstacles to communication. *Language Teaching* 42(4). 476–490. [doi](#) 10.1017/S026144480800551X
- Fernández, María del Carmen. 2013. El desarrollo de la conciencia fonética en edades tempranas. *Linred: Lingüística en la Red* 11. 1–25
- Freitas, Maria João, Celeste Rodrigues, Teresa Costa & Adelina Castelo. 2012. *Os sons que estão dentro das palavras: descrição e implicações para o ensino do Português como língua materna*. Colibri. [↗](#)
- Gil, Juana. 2007. *Fonética para profesores de español*. Arco
- Gomes, Aline. 2021. Produções fonéticas da vibrante múltipla /r/ por aprendentes baianos e paulistas de espanhol como língua estrangeira:

- analizando fatores extralinguísticos. *InterteXto* 14(1). 9–37. [doi 10.18554/it.v14i1.4723](https://doi.org/10.18554/it.v14i1.4723)
- João, Veloso, Ana Isabel Fernandes & Fátima Silva. 2024. *Dicionário fonético do português europeu*. Universidade do Porto. [↗](#)
- Levin, John. 2022. Teaching pronunciation: Truths and lies. Em Camilla Bardel, Christina Hedman, Katarina Rejman & Elisabeth Zetterholm (eds.), *Exploring Language Education: Global and local perspectives*, 39–72. Stockholm University Press
- Lipski, John. 2008. *Varieties of Spanish in the United States*. Georgetown University Press
- Listerri, Joaquim. 2003. La enseñanza de la pronunciación. *Cervantes* 4. 91–114
- Neto, José Rodrigues. 2014. Vogais espanholas: interferências da língua portuguesa e dificuldades de pronúncia. *Anais V SETEPE* 5. 1–16
- Neto, José Rodrigues & Maria Regina de Oliveira. 2023. O ensino de pronúncia no curso de letras-língua espanhola na UERN: desafios e estratégias. *Veredas – Revista de Estudos Linguísticos* 27(2). 1–20. [doi 10.34019/1982-2243.2023.v27.41486](https://doi.org/10.34019/1982-2243.2023.v27.41486)
- Odisho, Edward. 2007. A multisensory, multi-cognitive approach to teaching pronunciation. *Linguística - Revista de estudos linguísticos da Universidade do Porto* 2. 3–28. [↗](#)
- Perozzo, Reiner Vinicius & Felipe Flores Kupske. 2024. Percepção fônica não nativa e ensino de línguas: amalgamando fonologia e práticas pedagógicas. *Revista X* 19(3). 941–972. [↗](#)
- Polivanov, Evgueni. 1978. A percepção dos sons de uma língua estrangeira. Em *Círculo Lingüístico de Praga: estruturalismo e semiologia*, 113–128. Globo
- Rodríguez, Alfredo Macieira. 2004. Aspectos comparativos entre o português e o espanhol. Em *II Congresso Nacional de Lingüística e Filologia*, [↗](#)
- Silveira, Rosane & Tharen Teixeira Souza. 2011. A percepção e a produção das fricativas alveolares da língua portuguesa por hispano-falantes. *Revista de Estudos da Linguagem* 19(2). 167–184. [doi 10.17851/2237-2083.19.2.167-184](https://doi.org/10.17851/2237-2083.19.2.167-184)
- Stein, Cirineu Cecote. 2011. O conhecimento fonético acústico-articulatório e o ensino de língua estrangeira. *Signum*. *Estudos da Linguagem* 14(2). 355–374. [doi 10.5433/2237-4876.2011v14n2p355](https://doi.org/10.5433/2237-4876.2011v14n2p355)
- Telles, Luciana Pilatti & Jéssica Lung. 2024. Dia de luz, festa de sol: reflexões sobre a construção de uma unidade didática visando ao desenvolvimento de consciência fonológica da regra variável de alçamento de vogais médias-altas em um curso de pronúncia em português para falantes de outras línguas. *Letras de Hoje* 58(1). 1–14. [doi 10.15448/1984-7726.2023.1.44801](https://doi.org/10.15448/1984-7726.2023.1.44801)
- Veloso, João. 2012. Vogais centrais do português europeu contemporâneo: uma proposta de análise à luz da fonologia dos elementos. *Letras de Hoje* 47(3). 234–243. [↗](#)

Norma, ús i interferència: biaixos lingüístics en els models de llenguatge en català

Norm, use and interference: linguistic biases in Catalan language models

Mireia Almena Rodríguez ✉

Departament de Traducció i Ciències del Llenguatge, Universitat Pompeu Fabra

Thomas Brochhagen ✉ 

Departament de Traducció i Ciències del Llenguatge, Universitat Pompeu Fabra

Resum

Els Models de Llenguatge Extensos influeixen cada cop més en la comunicació escrita, fet que planteja reptes per a llengües minoritàries com el català. Aquest estudi quantifica els biaixos lingüístics de sis models causals de llenguatge en català, analitzant-ne les preferències per construccions gramaticals normatives enfront de les no normatives, especialment sota la influència potencial de la interferència del castellà. A partir d'un corpus propi de parelles mínimes, hem avaluat models tant monolingües com multilingües comparant les seves preferències per a cada variant (no) normativa. Els resultats indiquen que no hi ha una diferència entre els models monolingües i multilingües en la seva preferència per les construccions normatives. Tanmateix, la possible interferència del castellà redueix de manera notable la preferència per les formes normatives en tots els models analitzats. Aquesta reducció és més pronunciada en models multilingües, indicant més influència del castellà. Aquests resultats suggereixen que els biaixos dels models reflecteixen la prevalença d'usos no normatius en les seves dades d'entrenament, a causa d'influència del castellà. Això subratlla la importància d'avaluar aquestes tecnologies per a informar la política lingüística i comprendre'n l'impacte en l'evolució de la llengua.

Paraules clau

models de llenguatge; biaixos; norma; us; interferència; català

Abstract

Large Language Models are increasingly influencing written communication. This poses challenges for minority languages such as Catalan. This study quantifies the linguistic biases of six Catalan language models, analyzing their preferences for normative versus non-normative grammatical constructions, especially for cases where there can be interference from Spanish. Using a corpus of minimal pairs, we evaluate both monolingual and multilingual models by comparing their preferences for each (non-)normative variant. The results indicate that there is no differ-

ence between monolingual and multilingual models in their preference for normative constructions. However, cases where there can be interference from Spanish markedly reduce the preference for normative forms across all analyzed models. This reduction is more pronounced in multilingual models, indicating stronger influence from Spanish. These findings suggest that the models' biases reflect the prevalence of non-normative usage in their training data, due to influence from Spanish. This underscores the importance of evaluating these technologies to inform language policy and understand their impact on language evolution.

Keywords

language models; biases; norm; use; interference; catalan

1. Introducció

Les tecnologies del llenguatge tenen un paper cada vegada més important en la nostra vida quotidiana i en la manera com ens comuniquem. En particular, en pocs anys, s'ha estès l'ús de Models de Llenguatge Extensos generatius (LLMs, per les seves sigles en anglès). Ara per ara ja són eines que es fan servir per realitzar tasques rudimentàries, com redactar correus electrònics, o resumir i reestructurar textos. Podem suposar que, si l'adopció d'aquesta tecnologia segueix les tendències actuals, els LLMs formaran una part inextricable del pre- i post-processament del nostre ús de llenguatge escrit, com ho són ara els correctors automàtics d'ortografia i funcions d'emplenament automàtic.

Aquest desenvolupament té conseqüències per als llenguatges: si la proporció de text generada, o almenys modificada, per LLMs és cada vegada més gran (Liang et al., 2025), l'ús del llenguatge dels models —de biaixos a idiosincràsies estilístiques— no només retroalimentarà noves ge-



DOI: 10.21814/lm.18.1.497

This work is Licensed under a

Creative Commons Attribution 4.0 License

neracions de models, sinó que també pot afectar l'ús dels parlants. En altres paraules, els models no només reproduïen el llenguatge, sinó que també l'influencien. Això pot afectar més les llengües minoritàries, com el català, que són nodrides per un volum menor de contingut escrit —amb una esfera de parlants més petita— que altres llengües. En aquest estudi abordem aquesta temàtica, analitzant les preferències lingüístiques de models de llenguatge en català. D'una banda, volem quantificar els biaixos i les preferències per usos del llenguatge (no) normatius que tenen els models actuals. De l'altra, amb aquesta quantificació volem obrir un debat informat sobre el tipus de preferències lingüístiques que haurien de reproduir.

Un altre factor que pot afectar especialment les llengües minoritàries és l'ús generalitzat de models multilingües (Schut et al., 2025). Aquests models sovint s'han entrenat amb un volum més gran de dades —i sovint de més qualitat— en llengües com l'anglès, en comparació amb llengües com el català o el gallec. Aquest és el cas de tots els models multilingües catalans que analitzem aquí (vegeu §2). En conseqüència, les estructures gramaticals i el lèxic, entre altres, de les llengües amb més recursos poden afectar les llengües amb menys recursos. Diversos estudis ja han analitzat els biaixos lingüístics que presenten alguns models multilingües, especialment la seva tendència a afavorir estructures gramaticals i conceptuals angleses quan generen text en altres llengües (Papadimitriou et al., 2023; Wendler et al., 2024; Brinkmann et al., 2025; Schut et al., 2025). Però encara no sabem gaire sobre biaixos i interferències causades en aquests models entre altres llengües, com ho són el castellà i el català. En el cas del català, no queda clar fins a quin punt els models de llenguatge, tant monolingües com multilingües, fan servir estructures normatives, en comparació amb no-normatives. Amb aquest estudi, pretenem tancar aquesta bretxa.

Donat el fet que aquestes tecnologies acabaran influïnt en la nostra parla, la qual, al seu torn, alimentarà els nous models, considerem essencial rastrejar i avaluar les preferències lingüístiques dels models actuals; i analitzar fins a quin punt aquestes preferències estan influenciades per altres llengües. Això és el que ens proposem en aquest article: elaborem un corpus d'avaluació de parelles mínimes (no) normatives del català amb el qual avaluem sis models, tant mono- com multilingües.

2. Material i Mètodes

Tot el codi i el conjunt de dades que creem per avaluar models són disponibles a GitHub¹. Comencem amb una descripció dels models triats i després descrivim les dades per avaluar-los.

Models Hem seleccionat un conjunt de sis models, dos de monolingües i quatre de multilingües, capaços de generar i predir contingut en català. Amb models *multilingües* catalans ens referim a models que indiquen oficialment haver estat entrenats, entre d'altres, amb dades en català. Els models multilingües que estudiem són: Salamandra-7B² (Gonzalez-Agirre et al., 2025), un model base entrenat en 35 llengües europees (39% anglès, 16% castellà, 2% català); BLOOM-7.1B³ (Le Scao et al., 2023) un model base entrenat en 45 llengües (31% anglès, 11% castellà, 1% català); XGLM-7.5B⁴ (Lin et al., 2022), un model base entrenat en 31 llengües (49% anglès, 5% castellà, 0.4% català); i ALIA-40B⁵ (Gonzalez-Agirre et al., 2025), una variant més gran del model Salamandra, amb la mateixa cobertura lingüística.

Amb models *monolingües* catalans ens referim a models que han passat per un procés final d'aprenentatge que els ajusta específicament al català. Els models monolingües que estudiem són: CataLlama-8B v0.1⁶, basat en Llama-3 i entrenat durant un epoch amb 331 milions de tokens en català; i Aitana-6.3B⁷, un model basat en FLOR-6.3B i entrenat durant 2 epochs amb 1.3 milions de tokens en valencià. Excepte ALIA-40B i Salamandra-7B, tots aquests models han estat recentment avaluats pel que fa a les seves capacitats generals a IberoBench (Baucells et al., 2025). Però cap d'ells ha estat estudiat pel que fa a preferències de llenguatge (no) normatiu o possibles interferències.

La distinció que fem entre models monolingües i multilingües és habitual a la literatura. Cal esmentar, tanmateix, que la majoria dels models monolingües actuals han estat entrenats amb dades d'altres llengües, ja sigui en etapes inicials (abans del *fine tuning* per a la llengua objectiu), o bé de manera no intencionada, atès que les dades d'entrenament no s'han depurat fins al

¹<https://github.com/mireia-almena/NormaUsiInterferencia>

²<https://huggingface.co/BSC-LT/salamandra-7b>

³<https://huggingface.co/bigscience/bloom-7b1>

⁴<https://huggingface.co/facebook/xglm-7.5B>

⁵<https://huggingface.co/BSC-LT/ALIA-40b>

⁶<https://huggingface.co/catallama/CataLlama-v0.1-Base>

⁷<https://huggingface.co/gplsi/Aitana-6.3B>

punt de poder garantir que no s'hagin emprat dades d'una altra llengua en l'entrenament. Aquest és un fet a tenir en compte, particularment en avaluar models de llengües amb menys recursos. En conseqüència, i fet rellevant per al cas que ens ocupa, no és improbable que els models monolingües catalans hagin estat entrenats també, si més no parcialment, amb dades del castellà. Tornem a abordar aquesta qüestió en la discussió dels resultats.

Vam descartar dos altres models potencialment rellevants de l'avaluació: RoBERTa-ca i CataLlama-v0.2-Base. RoBERTa-ca⁸ és un masked language model, amb un objectiu d'entrenament diferent del de la resta de models causals que avaluem. Atès que la nostra mètrica d'avaluació es basa en la negative log-likelihood (NLL, vegeu més avall), vam descartar models no causals perquè la manera de calcular la NLL és diferent, i aquesta diferència podria introduir soroll a l'anàlisi. CataLlama-v0.2⁹, una fusió entre CataLlama-v0.1 i Meta-Llama-3-8B-Instruct, va ser descartat perquè vam considerar que la fusió podria fer el model “més multilingüe” que la resta de membres de la categoria monolingüe, preferint avaluar la seva versió 0.1 sense fusió amb un model d'instrucció en anglès. Una anàlisi post hoc revela que aquestes dues decisions de descart no afecten qualitativament els nostres resultats: incloue RoBERTa-ca i CataLlama-v0.2 —analitzant 4 models multilingües i 4 de monolingües— suggereix les mateixes tendències que els resultats obtinguts amb l'exclusió d'aquests dos models (vegeu l'apèndix). Per consegüent, els resultats reportats a la secció §3 són robustos davant aquesta decisió.

Corpus d'avaluació Hem construït un corpus d'avaluació amb vint oracions per cada un de vuit tipus d'estructures gramaticals en català que poden donar lloc a la producció de respostes no normatives. Quatre d'aquestes estructures són susceptibles a generar ús fora de la norma per interferència amb el castellà. Les altres quatre, no relacionades amb la interferència, s'han inclòs per observar si alguns models tracten de manera diferent les construccions influïdes per una altra llengua en comparació amb aquells que són simplement desviacions de la norma. Il·lustrem cada tipus d'estructura amb una oració exemplar, indicant com a primera opció la normativa i com a segona la no normativa:

1. Em refereixo $\{, a\}$ que va cantar.
2. He vist $\{el, al\}$ cotxe nou.
3. Al restaurant fa molta olor $\{de, a\}$ fregit.
4. Insisteixen $\{a, en\}$ tornar cap a casa.
5. No sóc gens propens $\{a, de\}$ enfadar-me per tonteries.
6. Van cancel·lar el vol $\{amb, en\}$ motiu del mal temps.
7. Cal confiar $\{en, amb\}$ el criteri dels experts.
8. Si no vas $\{amb, en\}$ compte, et pots fer mal.

Aquest corpus permet avaluar la qualitat lingüística dels models a partir de frases en què coexisteix una forma normativa i una altra de no normativa, sovint present en l'ús dels parlants. Com il·lustrat dalt, totes les estructures incloses en els conjunts de dades es basen en construccions gramaticals relacionades amb preposicions, ja que es tracta d'una categoria tancada que ofereix menys variabilitat en les respostes dels models. Això facilita la comparació de les respostes com a normatives o no normatives.

Per a cadascuna de les 160 oracions, calculem la negative log-likelihood (NLL) assignada a la seva versió (a) normativa i (b) no normativa per cada model. En altres paraules, per comparar dos oracions, calculem el logaritme negatiu de les seves probabilitats:

$$\begin{aligned} P(W) &= P(w_1, \dots, w_n) \\ &= P(w_1) \times P(w_1 \mid w_2) \times \dots \times \\ &\quad P(w_n \mid w_1, \dots, w_{n-1}) \end{aligned} \quad ,$$

per a una oració donada W , composta d'una seqüència de tokens de w_1 a w_n . En termes pràctics, extraïem la NLL mitjana d'una seqüència calculant la seva pèrdua (*loss*; vegeu ‘analysis.py’ al repositori per a detalls d'implementació). Com les parelles d'oracions que comparem només es diferencien al token que les diferencia entre normatives i no normatives, quantifiquem la preferència del model entre oració normativa (a) i oració no-normativa (b) com la diferència entre la NLL mitjana de l'oració (b) i la de l'oració (a). Un valor positiu indica una preferència per la versió normativa, i un valor negatiu, per la no normativa.

Cal esmentar que, mentre que la majoria de parelles d'oracions tenen la mateixa longitud, aquelles que impliquen elisió necessàriament difereixen per una paraula (p. ex., caigudes de preposició). Per assegurar-nos que els nostres resultats no estiguin influïts per possibles artefactes

⁸<https://huggingface.co/BSC-LT/RoBERTa-ca>

⁹<https://huggingface.co/catallama/CataLlama-v0.2-Base>

introduïts per longituds de seqüència diferenciats, també vam dur a terme l'anàlisi només amb el subconjunt de frases de longituds iguals. Trobem que els resultats que reportem a §3 s'obtenen independentment d'aquesta qüestió (vegeu l'apèndix per a estimacions numèriques quan només es consideren frases d'igual longitud).

Com que la magnitud de la diferència en NLL depèn del model —entre altres factors, està influïda per la mida del vocabulari— recodifiquem aquesta informació com un resultat binari de preferència a favor o en contra de (a) enfront de (b). Dit d'una altra manera, obtenim 160 judicis per model sobre si assignen una probabilitat més alta a la versió normativa o a la no normativa d'una oració. Abans de continuar a l'avaluació, expliquem breument cadascun dels vuit tipus d'estructures analitzades.

Estructura 1: Caiguda de preposició davant la construcció *que*. En registres formals, segons la GIEC (Institut d'Estudis Catalans, 2024, §26.4.1) s'evita el contacte de les preposicions *a*, *de*, *amb*, *en* i la conjunció *que*, de manera que s'elideixen aquestes preposicions.

Estructura 2: Preposició *a* davant d'objecte directe. Aquesta construcció és més comuna en castellà, especialment en complements directes animats. Segons la GIEC (Institut d'Estudis Catalans, 2024, §19.3.2), el complement directe en català no va precedit de cap preposició, excepte en casos d'ambigüitat amb el subjecte o l'alteració de l'ordre bàsic dels constituents oracionals, que no s'han inclòs en aquest corpus d'avaluació.

Estructura 3: Preposició *de* introductora de complements regits. Segons la normativa, s'han d'introduir amb la preposició *de*, i no pas *a*, certs complements regits, que per influència del castellà se solen construir malament (Murtra et al., 2025).¹⁰ Alguns exemples normatius són: *contradictori de*, *diferent de*, *fer cas de*, o *rebuig de*.

Estructura 4: Alternança de preposicions davant d'infinitiu. Els verbs o altres categories que tenen de complement o adjunt un sintagma nominal introduït per *en* o *amb*, aquestes preposicions alternen amb *a* o *de*, quan introdueixen una subordinada en infinitiu. Tanmateix, segons la GIEC (Institut d'Estudis Catalans, 2024, §26.5.2), la solució preferible en els registres formals és el canvi a *a* o *de*.

Estructura 5: Ús no normatiu de *propens de*. Aquesta estructura no és conseqüència directa del contacte amb el castellà, on el verb i la seva preposició són anàlegs a la normativa catalana. La forma normativa en català és *propens a* (Institut d'Estudis Catalans, 2007).

Estructura 6: Ús no normatiu de la locució *en motiu de*. Com l'estructura anterior, no és conseqüència del contacte amb el castellà. La forma normativa és *amb motiu de* (Serveis d'Assessorament Lingüístic del Parlament de Catalunya, 2007, Fitxa 7315/2).

Estructura 7: Usos no normatius de règims verbals no atribuïbles a la interferència del castellà: *confiar amb*, *continuar amb* i *tractar (de)*. La normativa, estableix *confiar en* (Institut d'Estudis Catalans, 2007), i considera *continuar* un verb transitiu que no admet *amb* (Institut d'Estudis Catalans, 2009). En el sentit de 'prendre alguna cosa com a objecte d'estudi o discussió', la normativa dicta que *tractar* ha de construir-se amb *de* (Institut d'Estudis Catalans, 2009).

Estructura 8: Ús no normatiu de la locució *anar en compte*. Com l'estructura anterior, no és conseqüència del contacte amb el castellà. La forma normativa és *anar amb compte* (Institut d'Estudis Catalans, 2024, §19.3.3.2.).

La llista completa d'oracions és disponible en format CSV al nostre repositori. També les produïm totes a l'apèndix.

3. Resultats

A continuació, analitzem la preferència dels models respecte a les nostres parelles mínimes de català (no) normatiu. En particular, quantifiquem si hi ha una diferència entre les preferències dels models monolingües i multilingües; i si les preferències respecte a l'ús normatiu varia en funció de si hi pot haver efectes d'interferència amb el castellà. Per a això, fem una regressió logística amb les 960 preferències obtingudes (160 per model). La regressió estima els efectes principals sobre les preferències per oracions (no) normatives en funció de: (i) la monolingüïtat vs. multilingüïtat de cada model; (ii) possible interferència amb el castellà (oracions del tipus Estructura 1–4 vs. Estructura 5–8); i (iii) la interacció entre aquests dos factors, amb interseccions estimades per oració i per model ("by-sentence" i "by-model" intercepts/efectes aleato-

¹⁰<https://www.uoc.edu/portal/ca/servei-linguistic/criteris/gramatica/preposicions/index.html>

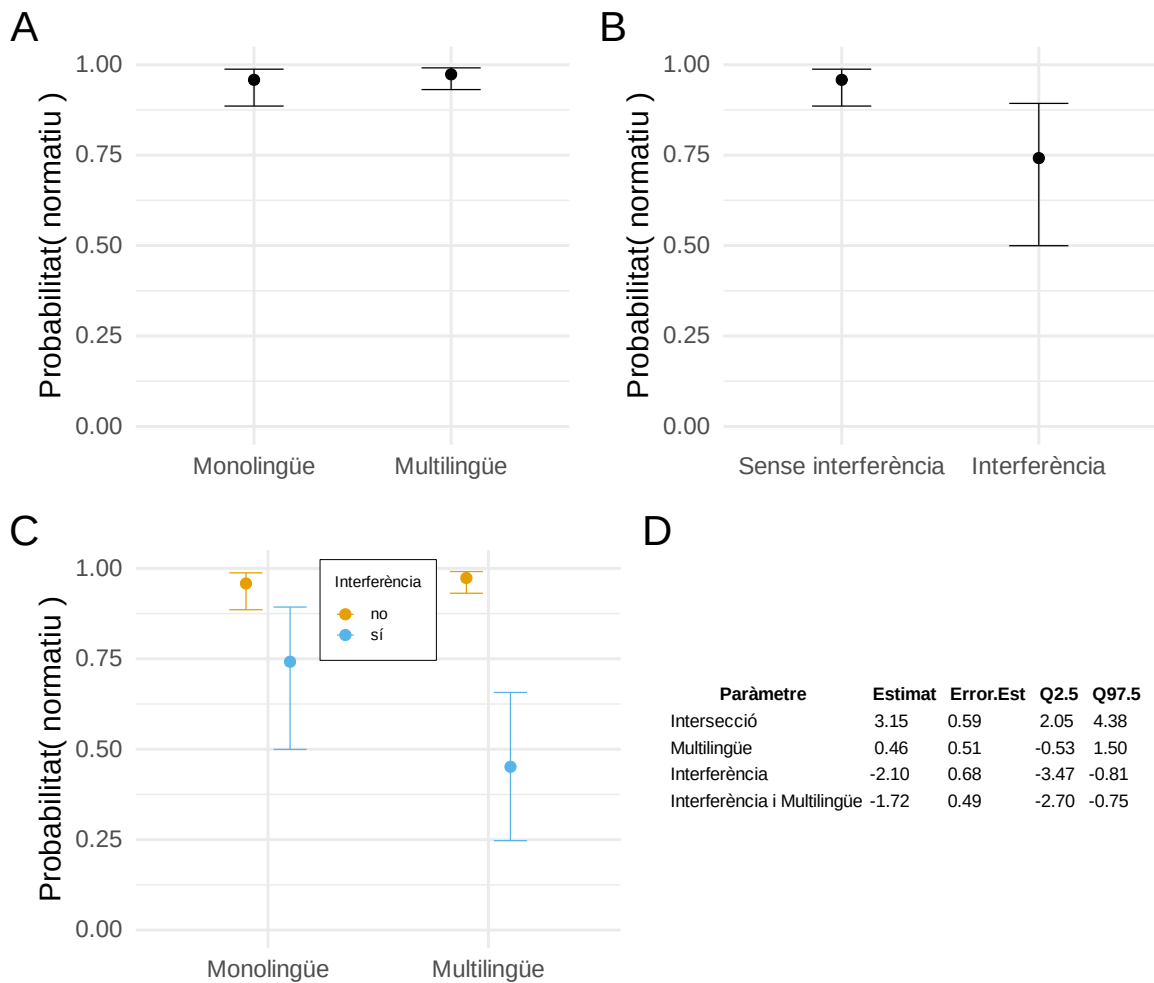


Figura 1: Estimacions del nostre model de regressió logística. A: efecte marginal de la monolingüïtat vs. multilingüïtat sobre la probabilitat de preferir una oració normativa davant la seva alternativa no normativa. B: efecte marginal de preferències entre oracions amb o sense possible interferència amb el castellà. C: interacció entre multilingüïtat i interferència. D: estimacions numèriques del model per a cada efecte amb ICs del 95%.

ris).¹¹ En altres paraules, estimem si el fet que un model sigui multilingüe, sense un ajust final específic per al català, afecta les preferències per les formes normatives; si la interferència amb el castellà les afecta; i, finalment, si hi ha un efecte específic d'interacció d'interferència en el cas dels models multilingües.

El model de regressió bayesià es va estimar fent servir el paquet de R *brms* (Bürkner, 2017) com a interfície amb Stan (Carpenter et al., 2017). Es va inspeccionar per descartar patologies en l'estimació. En particular, vam avaluar la fiabilitat de la mida de la mostra d'estimacions (més de 0,001 mostres efectives per transició) i vam comprovar que el valor de R^2

fragmentat fos inferior a 1,05, que la Bayesian Fraction of Missing Information fos superior a 0,2, i l'absència de trajectòries saturades. La validació creuada presenta valors de pareto- $k < 0,7$, cosa que suggereix estimacions fiables de la densitat predictiva logarítmica esperada (Vehtari et al., 2017).

La Figura 1 mostra els nostres resultats principals. Els tres primers panells mostren els efectes marginals dels tres factors que estudiem. El quart mostra el resum numèric sencer del model de regressió.

La Figura 1A mostra l'efecte marginal estimat de la multilingüïtat sobre la probabilitat de preferir una versió normativa d'una oració. Com s'il·lustra, i també es reflecteix en l'estimació numèrica (Figura 1D), no hi ha una diferència clara entre models mono- i multilingües, en si, pel que fa a la seva preferència per la versió normati-

¹¹En sintaxi de *lme4/brms*: $\text{preferencia} \sim 1 + \text{multiling} * \text{interferencia} + (1|\text{oracio}) + (1|\text{model})$ amb “preferència” codificada amb 1 si la diferència entre la NLL no-normativa i la normativa és positiva.

va. En contrast, la Figura 1B mostra l'efecte entre oracions amb interferència (Estructura 1–4) i sense interferència (Estructura 5–8) del castellà. Aquí veiem molta més variació, i el model de regressió suggereix que les oracions amb interferència tenen una probabilitat més baixa de comportar una preferència per la versió normativa. En altres paraules: El castellà interfereix amb el català normatiu. La Figura 1C mostra la interacció marginal entre la multilingüïtat i la interferència: Tot i que a la Figura 1B ja s'observa que els models mostren efectes d'interferència, aquest efecte és encara més pronunciat en els models multilingües. La Figura 1D ofereix un resum numèric de les estimacions de la regressió.

En conjunt, aquests resultats suggereixen que no hi ha una diferència clara entre els models monolingües i multilingües pel que fa a la seva preferència general per les oracions normatives davant alternatives no normatives (Figura 1A i segona fila de la Figura 1D). La interferència amb el castellà disminueix la preferència dels models per l'ús normatiu (Figura 1B i tercera fila de la Figura 1D), i hi ha una disminució addicional de la preferència per la normativitat en els models multilingües en les oracions amb interferència en comparació als monolingües (Figura 1C i quarta fila de la Figura 1D).

4. Discussió

Hi ha tres resultats principals en aquesta anàlisi. Primer, a grans trets, els models monolingües i multilingües no difereixen en la seva preferència per oracions normatives. Segon, quan hi ha la possibilitat d'interferència amb el castellà, la preferència de tots dos tipus de models es veu afectada: mostren una major inclinació cap a variants no normatives. Tercer, l'efecte de la interferència difereix entre models monolingües i multilingües: La preferència contra les variants normatives és molt més forta en els models multilingües.

Aquests resultats estan en línia amb les nostres expectatives. Trobem diferències entre models mono- i multilingües allà on són més probables: quan les dades d'altres llengües estretament relacionades, en aquest cas el castellà, apunten en una direcció diferent. No obstant això, és notable que no trobem una diferència per als casos sense interferència. Això vol dir que, per aquests casos i en un domini tancat com les preposicions, els models multilingües desenvolupen competències comparables a les dels models monolingües, independentment que les dades amb què s'han entrenat estiguin desequilibrades (cf. Brinkmann et al., 2025).

Això ens remet a la discussió de l'apartat 2 relativa a la distinció terminològica entre models monolingües i multilingües. Si més no pel que fa a les dades aquí analitzades, constatem que la distinció té rellevància pràctica, però que la seva rellevància s'ha de jutjar cas per cas. També es podria haver esperat que els models multilingües tinguessin preferències més febles per la norma en els casos sense interferència. Aquest resultat és positiu: els models actualment més populars amb usuaris no mostren desavantatges en almenys alguns aspectes en relació amb models especialitzats en català, que moltes vegades no arriben fins als usuaris.

El segon i tercer resultat —un efecte d'interferència per a tots els models, i particularment per als multilingües— convida a la reflexió i, potser, a l'acció. D'una banda, la variació en dades i entrenament no sembla alterar el fet que la preferència per la normativa disminueixi en contextos d'interferència amb el castellà. Això suggereix que les variants no normatives amb interferència que hem estudiat ja formen part de l'ús del català al qual s'exposen els models. El nostre objectiu aquí no és adjudicar si això hauria de comportar un canvi en la normativa; o bé si caldria reforçar l'ensenyament del català en aquells punts on divergeix del castellà; o bé si convindria preprocessar les dades amb què s'entrenen els models per evitar que l'ús no normatiu sigui propagat pels LLMs. El que sí que mostra clarament aquest resultat és la importància de verificar les preferències lingüístiques dels LLMs, especialment en llengües minoritàries, per tal de tenir aquesta discussió de manera informada; per elaborar polítiques lingüístiques que tinguin en compte l'impacte de les noves tecnologies; i per actuar en conseqüència.

De l'altre, la preferència encara més feble per la norma dels models multilingües demostra com la tecnologia podria actuar com una força de canvi lingüístic: amplificant patrons d'ús i propagant-los. Més en general, aquests resultats mostren que el desenvolupament de tecnologies lingüístiques adaptades a llengües específiques és important i s'hauria de dur a terme conjuntament amb les multilingües. El que funciona per a un conjunt de llengües pot no generalitzar-se a d'altres amb realitats tipològiques i sociolingüístiques diferents. Això és important tant per millorar els models mateixos (Baucells et al., 2025; Pérez-Mayos et al., 2021; Táboas García et al., 2025) com per detectar conseqüències posteriors potencialment inesperades com les que destaquem aquí.

Una anàlisi complementària que cal fer per posar l'ús lingüístic dels models en context és l'estudi dels parlants de català, cosa que deixem per a futures investigacions. Vam construir aquest corpus d'oracions precisament per estudiar no una diferència entre el que és “correcte” vs. “incorrecte”, sinó entre diferents usos. És ben sabut que molts parlants —catalans o d'altres— no segueixen necessàriament la norma en l'ús que fan de la llengua. Per tant, és probable que fins a cert punt l'efecte d'interferència que detectem és un reflex genuí del canvi lingüístic. Creiem que combinar aquest tipus d'estudi amb experiments conductuals amb subjectes humans és una línia molt prometedora per a futures recerques.

Una altra via prometedora per a futures investigacions és una anàlisi de les dades en brut amb què s'entrenen els models. Aquí ens hem centrat en el comportament lingüístic dels models després del seu entrenament. Però, per deslligar causalment què provoca les seves similituds i diferències, caldria tornar a les dades d'entrada i comparar les taxes de prevalença de les construccions avaluades. Això també aclariria fins a quin punt els LLM amplifiquen o simplement reproduïxen els patrons que observen durant l'entrenament; i com interactua aquest fet amb la prevalença de dades d'entrenament mono- vs. multilingües.

5. Conclusió

Aquest estudi s'ha centrat en l'anàlisi dels biaixos lingüístics dels models de llenguatge en català, avaluant-ne les preferències per construccions normatives enfront de les no normatives, i examinant la influència de la interferència lingüística del castellà. L'avaluació de sis models de llenguatge, tant monolingües com multilingües, a partir d'un corpus de parelles mínimes, revela dos resultats principals.

Primerament, no hem trobat una diferència generalitzada entre els models monolingües i els models multilingües. En segon lloc, la presència de possible interferència amb el castellà redueix la preferència per les variants normatives en tots els models analitzats. Aquesta disminució és marcadament més forta en els models multilingües, la qual cosa podria conduir a una major ampliació d'aquestes preferències en parla catalana generada o modificada per aquests models, que són molt més populars amb usuaris.

Aquests resultats també suggereixen que les variants no normatives influïdes pel castellà estan prou esteses en l'ús real de la llengua com per ser reflectides de manera consistent en les da-

des d'entrenament de tots els models. Donat el paper cada vegada més influent que les tecnologies del llenguatge tenen en la nostra comunicació quotidiana, la comprensió d'aquests biaixos és fonamental. Els models no només reproduïxen el llenguatge, sinó que també tenen el potencial d'influir-lo i propagar certs usos. Per tant, aquest estudi subratlla la importància de verificar les preferències lingüístiques dels LLMs per poder fomentar un debat informat.

Agraïments

Agraïm a Iria de Dios Flores, Antoni Oliver González i Lluís Padró pels seus comentaris i suggeriments. TB és finançat pel Ministerio de Ciencia, Innovación, y Universidades, l'Agencia Estatal de Investigación i el Fons Social Europeu Plus (RYC2023-045215-I MCIU/AEI/10.13039/501100011033), com també pel projecte EVOSIG PID2024-162668NA-I00 finançat pel MICIU/AEI/10.13039/501100011033/ FEDER, UE.

Referències

- Baucells, Irene, Javier Aula-Blasco, Iria de Dios-Flores, Silvia Paniagua Suárez, Naiara Perez, Anna Salles, Susana Sotelo Docio, Júlia Falcão, Jose Javier Saiz, Robiert Sepulveda Torres, Jeremy Barnes, Pablo Gamallo, Aitor Gonzalez-Agirre, German Rigau & Marta Villegas. 2025. IberoBench: A benchmark for LLM evaluation in Iberian languages. En *31st International Conference on Computational Linguistics (COLING)*, 10491–10519. [↗](#)
- Brinkmann, Jannik, Chris Wendler, Christian Bartelt & Aaron Mueller. 2025. Large language models share representations of latent grammatical concepts across typologically diverse languages. En *Conference of the Nations of the Americas Chapter of the Association for Computational Linguistics (NAACL)*, 6131–6150. [doi](#) 10.18653/v1/2025.naacl-long.312
- Bürkner, Paul-Christian. 2017. brms: An R Package for Bayesian Multilevel Models using Stan. *Journal of Statistical Software* 80(1). [doi](#) 10.18637/jss.v080.i01
- Carpenter, Bob, Andrew Gelman, Matthew D. Hoffman, Daniel Lee, Ben Goodrich, Michael Betancourt, Marcus Brubaker, Jiqiang Guo, Peter Li & Allen Riddell. 2017. Stan: A probabilistic programming language. *Journal of Statistical Software* 76(1). [doi](#) 10.18637/jss.v076.i01

- Gonzalez-Agirre, Aitor, Marc Pàmies, Joan Llop, Irene Baucells, Severino Da Dalt, Daniel Tamayo, José Javier Saiz, Ferran Espuña, Jaume Prats, Javier Aula-Blasco, Mario Mina, Iñigo Pikabea, Adrián Rubio, Alexander Shvets, Anna Sallés, Iñaki Lacunza, Jorge Palomar, Júlia Falcão, Lucía Tormo, Luis Vasquez-Reina, Montserrat Marimon, Oriol Pareras, Valle Ruiz-Fernández & Marta Villegas. 2025. Salamandra technical report. arXiv [cs.CL]. [doi 10.48550/arXiv.2502.08489](https://doi.org/10.48550/arXiv.2502.08489)
- Institut d'Estudis Catalans. 2007. Diccionari de la llengua catalana [diec2, online version]. [↗](#)
- Institut d'Estudis Catalans. 2009. Manual d'estil. la redacció i l'edició de textos (4a ed.). [↗](#)
- Institut d'Estudis Catalans. 2024. *Gramàtica de la llengua catalana [GIEC, online version]*. Barcelona: IEC. [doi 10.2436/10.2500.08.1](https://doi.org/10.2436/10.2500.08.1)
- Le Scao, Teven, Angela Fan, Christopher Akiki, Ellie Pavlick & et al. 2023. BLOOM: A 176b-parameter open-access multilingual language model. arXiv [cs.CL]. [doi 10.48550/arXiv.2211.05100](https://doi.org/10.48550/arXiv.2211.05100)
- Liang, Weixin, Yaohui Zhang, Zhengxuan Wu, Haley Lepp, Wenlong Ji, Xuandong Zhao, Hancheng Cao, Sheng Liu, Siyu He, Zhi Huang, Diyi Yang, Christopher Potts, Christopher D. Manning & James Zou. 2025. Quantifying large language model usage in scientific papers. *Nature Human Behaviour* [doi 10.1038/s41562-025-02273-8](https://doi.org/10.1038/s41562-025-02273-8)
- Lin, Xi Victoria, Todor Mihaylov, Mikel Artetxe, Tianlu Wang, Shuohui Chen, Daniel Simig, Myle Ott, Naman Goyal, Shruti Bhosale, Jingfei Du, Ramakanth Pasunuru, Sam Shleifer, Punit Singh Koura, Vishrav Chaudhary, Brian O'Horo, Jeff Wang, Luke Zettlemoyer, Zornitza Kozareva, Mona Diab, Veselin Stoyanov & Xian Li. 2022. Few-shot learning with multilingual language models. arXiv [cs.CL/cs.AI]. [doi 10.48550/arXiv.2112.10668](https://doi.org/10.48550/arXiv.2112.10668)
- Murtra, Pilar, Thomas Bell, David Cullen, Jordi Gavaldà, Cristina López, Xavier Marzal & Alba Pérez. 2025. Servei lingüístic de la Universitat Oberta de Catalunya. [↗](#)
- Papadimitriou, Isabel, Kezia Lopez & Dan Jurafsky. 2023. Multilingual BERT has an accent: Evaluating English influences on fluency in multilingual models. En *Findings of the Association for Computational Linguistics*, 1194–1200. [doi 10.18653/v1/2023.findings-eacl.89](https://doi.org/10.18653/v1/2023.findings-eacl.89)
- Pérez-Mayos, Laura, Alba Táboas García, Simon Mille & Leo Wanner. 2021. Assessing the syntactic capabilities of transformer-based multilingual language models. En *Findings of the Association for Computational Linguistics*, 3799–3812. [doi 10.18653/v1/2021.findings-acl.333](https://doi.org/10.18653/v1/2021.findings-acl.333)
- Schut, Lisa, Yarin Gal & Sebastian Farquhar. 2025. Do multilingual LLMs think in English? En *ICLR 2025 Workshop on Building Trust in Language Models and Applications*, [↗](#)
- Serveis d'Assessorament Lingüístic del Parlament de Catalunya. 2007. Optimot, consultes lingüístiques. [↗](#)
- Táboas García, Alba, Piotr Przybyła & Leo Wanner. 2025. Exploring morphology-aware tokenization: A case study on Spanish language modeling. En *Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 30493–30506. [doi 10.18653/v1/2025.emnlp-main.1552](https://doi.org/10.18653/v1/2025.emnlp-main.1552)
- Vehtari, Aki, Andrew Gelman & Jonah Gabry. 2017. Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing* 27. 1413–1432. [doi 10.1007/s11222-016-9696-4](https://doi.org/10.1007/s11222-016-9696-4)
- Wendler, Chris, Veniamin Veselovsky, Giovanni Monea & Robert West. 2024. Do Llamas work in English? on the latent language of multilingual transformers. En *62nd Annual Meeting of the Association for Computational Linguistics*, 15366–15394. [doi 10.18653/v1/2024.acl-long.820](https://doi.org/10.18653/v1/2024.acl-long.820)

Apèndix

Comprovacions de robustesa La Taula 1 mostra les estimacions d'un model de regressió que inclou CataLlama-v0.2 i RoBERTa-ca. L'estructura del model de regressió i tots els altres detalls són els mateixos que en el text principal. Per a RoBERTa-ca, la (pseudo-)negative log likelihood es calcula prenent el logaritme del likelihood de les oracions amb les preposició clau emascarades.

Com es pot observar, les estimacions són semblants a les obtingudes sense aquests dos models (comparar amb la Figura 1D). Principalment, l'efecte d'interferència és més gran amb la inclusió d'aquests dos models. Això suggereix que els resultats són relativament robustos davant de la variació en el conjunt de models estudiats, almenys pel que fa als monolingües.

Paràmetre	Estimat	Error.Est	Q2.5	Q97.5
Intersecció	3.45	0.51	2.51	4.56
Multilingüe	0.16	0.42	-0.64	0.96
Interferència	-2.72	0.60	-3.99	-1.62
Interferència i Multilingüe	-0.99	0.40	-1.82	-0.22

Taula 1: Estimacions numèriques per a un model de regressió que inclou CataLlama i RoBERTa-ca.

Paràmetre	Estimat	Error.Est	Q2.5	Q97.5
Intersecció	3.05	0.55	2.03	4.16
Multilingüe	0.69	0.56	-0.38	1.84
Interferència	-1.40	0.57	-2.55	-0.33
Interferència i Multilingüe	-1.59	0.52	-2.64	-0.59

Taula 2: Estimacions numèriques per a un model de regressió amb només estímuls amb el mateix nombre de paraules per parella (no) normativa.

La Taula 2 mostra les estimacions d'un model de regressió quan només es fan servir oracions d'igual longitud. S'obtenen les mateixes tendències que a la Fig. 1, cosa que indica que les longituds de seqüència diferencials ocasionals (p. ex., caigudes de preposició) no afecten els nostres resultats de maneres inesperades.

Estímuls A continuació es mostren els parells de oracions que vam analitzar. Com en els exemples del text principal, l'opció normativa és la primera dins dels claudàtors, abans de la coma. Les entrades buides indiquen que la respectiva variant (no) normativa deixa aquella posició buida.

Estructura 1

1. No m'havia adonat {, de} que estava mal col·locat.
2. Abans {, de} que comencés, ja ho sabíem.
3. Depèn {, de} que decideixin els experts.
4. No es recorda {, de} que li vaig explicar.
5. No m'havia adonat {, de} que estava equivocat.
6. Depèn {, de} que vulgin fer.
7. No es recorda {, de} que va veure aquell dia.
8. Tinc la certesa {, de} que arribaràs a temps.
9. M'alegro {, de} que ho sàpigues.
10. Estic convençut {, de} que tenim raó.
11. Aspirava {, a} que el triessin.
12. Contribuirà {, a} que el projecte tiri endavant.
13. S'oposen {, a} que la normativa es modifiqui.
14. Em refereixo {, a} que va fer una nova pel·lícula.

15. Em refereixo {, a} que va cantar.
16. Estic acostumat {, a} que m'avisin amb temps.
17. L'autoritat obliga {, a} que es pagui un peatge.
18. N'hi ha prou {, amb} que diguis la veritat.
19. Confio {, en} que tot sortirà bé.
20. Insisteix {, en} que tenim una altra opció.

Estructura 2

1. Això afecta {el, al} pintor solidari.
2. He vist {el, al} senyor calb.
3. L'entrenador felicitarà {els, als} jugadors després del partit.
4. L'empresa vol formar {els, als} seus treballadors.
5. Estic esperant {el, al} professor per resoldre un dubte.
6. Han acomiadat {el, al} empleats més antics de l'empresa.
7. Sempre ajuda {els, als} seus companys quan ho necessiten.
8. El professor ha avisat {els, als} alumnes sobre el canvi d'horari.
9. El jutge ha interrogat {els, als} testimonis del cas.
10. Han seleccionat {els, als} millors candidats per al lloc de feina.
11. La policia busca {el, al} culpable.
12. El jutge va absoldre {els, als} acusats.
13. Necessito trobar {el, al} metge.
14. El soroll va despertar {el, al} veí.
15. Vam citar {el, al} director de l'escola.
16. La notícia ha afectat {el, al} jove empresari.
17. Necessiten substituir {el, al} tècnic que ha marxat.
18. El llibre va impressionar {els, als} lectors més joves.
19. L'entrenador ha sancionat {el, al} jugador titular.
20. Van ajudar {els, als} excursionistes perduts.

Estructura 3

1. Al restaurant fa molta olor {de, a} fregit.
2. Tinc por {de, a} les aranyes.
3. El meu horari és diferent {del, al} teu.
4. La solució {del, al} problema és molt senzilla.
5. La noia va aprofitar per treure profit {de, a} la situació i va guanyar molts diners.
6. Va fer cas {de, a} les recomanacions del metge.
7. Sento una forta pudor {de, a} cremat a la cuina.
8. L'article fa esment {de, a} diverses investigacions recents.
9. Va mostrar un clar rebuig {de, a} la proposta.
10. Hauries de fer un bon repàs {de, a} la matèria abans de l'examen.

11. La seva versió és contradictòria {de, a} la que vam escoltar.
12. La pel·lícula és diferent {de, a} la novel·la.
13. No facis cas {de, a} les males notícies.
14. El presentador va fer esment {de, a} la polèmica.
15. Sentia una forta olor {de, a} tabac apagat.
16. Els nens tenien por {dels, als} gossos grans.
17. La noia porta un vestit {de, a} ratlles blanques i negres.
18. El jurat fa una menció especial {de, a} l'obra d'aquesta estudiant.
19. Això no implica cap increment {de, a} les {de, a}speses {de, a}l web.
20. Sempre ha volgut estudiar una carrera diferent {de, a} la que ha fet.

Estructura 4

1. Tenim molt d'interès {a, en} complir els nostres compromisos.
2. La idea consisteix {a, en} analitzar els models.
3. La majoria d'estudiants es van mostrar d'acord {a, en} fer una pausa.
4. Comptava {a, amb} fer això.
5. Insisteixen {a, en} tornar cap a casa.
6. Va mostrar interès {a, en} aprendre nous mètodes per fer aquests exercicis.
7. L'exercici es basa {a, en} explicar les circumstàncies del conflicte.
8. Ens complaem {a, en} acceptar la vostra col·laboració.
9. Engrescat {a, en} fundar la coral, tot el temps que hi dedica li sembla poc.
10. Estava tan capficat {a, en} completar la jugada, que no es va adonar que el temps s'havia acabat.
11. Tinc interès {a, en} aprendre més llengües.
12. El treball consisteix {a, en} establir una relació entre les dues variables.
13. Compte {a, amb} demanar el mateix a tots dos llocs
14. L'empresa es va comprometre {a, amb} lliurar la mercaderia abans de la data.
15. La subsistència es basava {a, en} intercanviar productes.
16. S'ha entossudit {a, en} comprar folis reciclats.
17. La democràcia consisteix {a, en} fer que el poble participi en la política.
18. No s'ha de capficar {a, en} recordar totes les dades de memòria.
19. Cal invertir un gran esforç {a, en} millorar l'atenció al client.
20. Ell es complau {a, en} resoldre trencaclosques complicats.

Estructura 5

1. És molt propens {a, de} agafar refredats a l'hivern.
2. Aquesta planta és propensa {a, de} les plagues si no la cuides bé.
3. Els adolescents són més propensos {a, de} prendre riscos innecessaris.
4. Les persones amb pells clares són propenses {a, de} les cremades solars.
5. El meu oncle és propens {a, de} explicar sempre les mateixes anècdotes.
6. Aquesta zona és propensa {a, de} les inundacions quan plou fort.
7. Som propensos {a, de} oblidar les claus a casa.
8. No sóc gens propens {a, de} enfadar-me per tonteries.
9. Les economies inestables són propenses {a, de} patir crisis financeres.
10. Aquest material és propens {a, de} l'oxidació si s'exposa a la humitat.
11. El meu veí és molt propens {a, de} exagerar les històries.
12. La pell sensible és propensa {a, de} reaccions al·lèrgiques.
13. Aquella noia és propensa {a, de} dir mentides sense motiu.
14. Els cotxes vells són propensos {a, de} avaries inesperades.
15. És propens {a, de} cometre errors quan està sota pressió.
16. Ets massa propens {a, de} canviar d'opinió a l'últim moment.
17. Aquesta màquina és propensa {a, de} errors si no es calibra bé.
18. El meu soci és propens {a, de} prendre decisions arriscades.
19. La nostra regió és propensa {a, de} grans nevades a l'hivern.
20. Les persones nervioses són propenses {a, de} mossegar-se les ungles.

Estructura 6

1. {Amb, En} motiu de la seva jubilació, li van organitzar una festa sorpresa.
2. La botiga va tancar {amb, en} motiu de les obres al carrer.
3. Es va organitzar un concert benèfic {amb, en} motiu de la recaptació de fons.
4. {Amb, En} motiu del Dia del Llibre, les llibreries fan descomptes.
5. Van cancel·lar el vol {amb, en} motiu del mal temps.
6. L'ajuntament va emetre un comunicat {amb, en} motiu de la nova normativa.

7. Es va tallar el trànsit {amb, en} motiu de la manifestació.
8. La reunió es va ajornar {amb, en} motiu de l'absència del director.
9. {Amb, En} motiu de la celebració, el museu ofereix entrada gratuïta.
10. L'empresa va fer una donació {amb, en} motiu de la campanya solidària.
11. {Amb, En} motiu de les festes, l'ajuntament va tancar.
12. La trobada es va convocar {amb, en} motiu de la crisi energètica.
13. Es va avançar l'hora {amb, en} motiu de la celebració del partit.
14. L'excursió es va cancel·lar {amb, en} motiu de la forta ventada.
15. Vam assistir a la conferència {amb, en} motiu de la presentació del llibre.
16. Han declarat dia festiu {amb, en} motiu del patronatge del poble
17. Vam fer una excepció {amb, en} motiu de la seva situació particular.
18. Vam llançar una edició especial {amb, en} motiu del centenari de la marca.
19. La subvenció va ser concedida {amb, en} motiu de les bones previsions.
20. {Amb, En} motiu dels preparatius, la pista d'esquí roman tancada.

Estructura 7

1. Pots confiar {en, amb} mi per guardar el secret.
2. No confio {en, amb} les seves promeses, sempre les incompleix.
3. Cal confiar {en, amb} el criteri dels experts.
4. He après a no confiar {en, amb} qualsevol que sembli amable.
5. Perquè la relació funcioni, heu de confiar l'un {en, amb} l'altre.
6. Confiem {en, amb} que tot sortirà bé al final.
7. És difícil confiar {en, amb} algú que t'ha mentit abans.
8. Els ciutadans confien {en, amb} la justícia del seu país.
9. Si us plau, continueu {, amb} la lectura del proper capítol.
10. El govern continuarà {, amb} les negociacions la setmana vinent.
11. Després d'aquesta interrupció continuem {, amb} el programa.
12. Malgrat tot, l'equip continua {, amb} aquest projecte.
13. Podem continuar {, amb} la feina demà al matí.

14. Continuem {, amb} la tasca que vam deixar a mig fer.
15. El debat tractarà {de, } les noves propostes de llei.
16. En aquest llibre l'autor tracta molt sintèticament {de, } la morfologia.
17. El més interessant del llibre és com tracta {de, } la sintaxi.
18. L'informe no tracta {de, } les qüestions més importants.
19. La pel·lícula tracta {de, } la vida de Mozart.
20. La conferència tractava {de, } la intel·ligència artificial.

Estructura 8

1. Vés {amb, en} compte quan creuis el carrer.
2. Aneu {amb, en} compte amb el terra, que rellisca.
3. Cal anar {amb, en} compte a l'hora de signar contractes.
4. Sempre vaig {amb, en} compte amb el que dic en públic.
5. Van anar {amb, en} molt de compte per no fer soroll.
6. Si no vas {amb, en} compte, et pots fer mal.
7. Anem {amb, en} compte, que no sabem què ens trobarem.
8. El metge li va dir que anés {amb, en} compte amb els greixos.
9. És una situació delicada, hem d'anar {amb, en} compte.
10. Aneu {amb, en} compte amb les estafes per internet.
11. El metge li va dir que anés {amb, en} compte amb la dieta.
12. Sempre va {amb, en} compte a l'hora d'invertir els seus diners.
13. L'advocat els va aconsellar anar {amb, en} compte amb les clàusules del contracte.
14. Ves {amb, en} compte quan entris al bosc, podria haver-hi animals.
15. El cuiner va haver d'anar {amb, en} compte amb l'oli calent.
16. Ves {amb, en} compte amb qui comparteixes la informació personal.
17. Si vas {amb, en} compte, evitaràs la majoria dels problemes.
18. Ella sempre va {amb, en} compte amb el que escriu a les xarxes socials.
19. El doctor va advertir-li que anés {amb, en} compte amb l'estrès.
20. Cal anar {amb, en} compte en les declaracions públiques.

Um Comitê de Classificadores para Anotação Automática de Linguagem Tóxica em Português sob Escassez de Dados

An Ensemble of Classifiers for Automatic Annotation of Toxic Language in Portuguese under Data Scarcity

Francisco Assis Ricarte Neto  
Instituto Federal do Piauí

Rafael Torres Anchiêta
Instituto Federal do Maranhão
 

Raimundo Santos Moura  
Universidade Federal do Piauí

Pedro de Alcântara dos Santos Neto  
Universidade Federal do Piauí

André Macedo Santana 
Universidade Federal do Piauí

Resumo

Mensagens com linguagem tóxica configuram um problema recorrente nas redes sociais, evidenciando a necessidade urgente de métodos automáticos eficazes para mitigar seu impacto. Em geral, tais abordagens dependem de grandes volumes de dados rotulados, cuja construção é onerosa e demorada, além de demandar considerável esforço humano no processo de anotação. Para enfrentar esse desafio, este trabalho propõe um comitê de classificadores voltado à anotação automática de linguagem tóxica em português, concebido para operar com um número reduzido de dados previamente anotados. O comitê integra três estratégias complementares: um método semi-supervisionado baseado em grafos heterogêneos, uma abordagem de aprendizagem *few-shot* e um método fundamentado em *Retrieval-Augmented Generation*, ambos baseados em grandes modelos de linguagem. A proposta é avaliada em múltiplos *corpora*, considerando conjuntos originais e filtrados por concordância total entre anotadores. Os resultados indicam que o comitê apresenta desempenho competitivo, superando o melhor método individual em até 2% em cenários de maior equilíbrio entre os classificadores constituintes e mantendo desempenho comparável nos demais, além de preservar concordância moderada a substancial com os rótulos originais, evidenciando seu potencial para a construção de recursos linguísticos anotados sob escassez de dados.

Palavras chave

linguagem tóxica; anotação automática; comitê de classificadores

Abstract

Messages containing toxic language are a recurring problem on social media, highlighting the urgent need for effective automatic methods to mitigate

their impact. Most existing approaches rely on large volumes of annotated data, which are costly, time-consuming, and highly labor-intensive. To address this challenge, this work proposes an ensemble of classifiers for the automatic annotation of toxic language in Portuguese, designed to operate under limited labeled data. The ensemble integrates three complementary strategies: a semi-supervised method based on heterogeneous graphs, a *few-shot* learning approach, and a *Retrieval-Augmented Generation* method, both grounded in large language models. The proposal is evaluated across multiple *corpora*, considering both their original versions and subsets filtered by total inter-annotator agreement. The results indicate that the ensemble exhibits competitive performance, surpassing the best individual method by up to 2% in scenarios of greater balance among the constituent classifiers and maintaining comparable performance in the remaining ones, while preserving moderate to substantial agreement with the original labels, demonstrating its potential for constructing annotated linguistic resources under data scarcity.

Keywords

toxic language; automatic annotation; ensemble of classifiers

1. Introdução

As redes sociais constituem poderosas ferramentas virtuais de interação humana, conectando pessoas em escala global e viabilizando a troca de informações, a expressão de opiniões e o debate de ideias. Entretanto, esses ambientes também têm favorecido a disseminação de diferentes formas de linguagem tóxica, por meio de mensagens que extrapolam os limites da liberdade de expressão e comprometem a quali-

dade da interação social. Esse tipo de conteúdo caracteriza-se pelo uso de linguagem inadequada, incluindo expressões profanas, insultos, ameaças e declarações ofensivas dirigidas a indivíduos ou grupos (Founta et al., 2018). Além das manifestações explícitas, a toxicidade pode ocorrer de forma mais sutil, por meio de comportamentos como comentários rudes, observações desrespeitosas ou ataques implícitos, frequentemente associados ao discurso de ódio ou a outras práticas comunicativas abusivas (Cjadams et al., 2017).

A relevância social da toxicidade em ambientes digitais tem sido amplamente discutida na literatura recente (Fortuna & Nunes, 2018; Poletto et al., 2020; Jahan & Oussalah, 2023). Embora existam marcos legais em diversos países que a tipificam como crime, como ocorre no Brasil (Brasil, 1989), além de políticas de moderação que preveem a remoção de conteúdos e, em casos mais graves, a exclusão de usuários das plataformas (X, 2026; Youtube, 2026), a linguagem tóxica permanece um problema de difícil contenção nos ambientes digitais. Diante do crescimento contínuo do número de usuários e do volume massivo de mensagens que exigem monitoramento, torna-se imprescindível o desenvolvimento de soluções automáticas para a detecção dessa linguagem. Do ponto de vista computacional, o problema tem sido tradicionalmente abordado como uma tarefa de classificação de texto, baseada majoritariamente em técnicas de Aprendizado de Máquina Supervisionado e de Aprendizagem Profunda (Jahan & Oussalah, 2023; Badjatiya et al., 2017). Tais abordagens pressupõem a disponibilidade de grandes volumes de dados anotados, o que constitui um problema, uma vez que os maiores esforços no desenvolvimento de recursos linguísticos são direcionados à língua inglesa (Poletto et al., 2020), enquanto línguas secundárias, como o português, permanecem sub-representadas. Embora alguns *corpora* em português tenham sido propostos (de Pelle & Moreira, 2017; Nascimento et al., 2019; Fortuna et al., 2019; Vargas et al., 2022; Leite et al., 2020; Trajano et al., 2024), a disponibilidade de conjuntos de dados anotados em larga escala ainda é limitada.

A construção manual de *corpora* é um processo oneroso e demorado, agravado, no caso da linguagem tóxica, pela subjetividade interpretativa inerente à tarefa, que frequentemente requer o reconhecimento de manifestações implícitas por meio de recursos pragmáticos da linguagem, bem como por influências socioculturais e ideológicas dos próprios anotadores, aspectos que podem comprometer a consistência e a confiabilidade dos

rótulos atribuídos (Aroyo et al., 2019; Hettiachchi et al., 2023). Adicionalmente, a anotação manual de textos tóxicos envolve impactos humanos significativos, uma vez que a exposição contínua a conteúdos violentos, abusivos ou perturbadores pode acarretar prejuízos psicológicos aos avaliadores¹.

Em resposta a essas limitações, a anotação automática de dados textuais tem emergido como uma alternativa promissora para a construção escalável de recursos linguísticos (Wang et al., 2021). Nesse contexto, abordagens semi-supervisionadas (Santos et al., 2022) e estratégias de *pseudo-labeling* (Dirting et al., 2022) vêm sendo exploradas para a construção e ampliação automática de *corpora* voltados à detecção de toxicidade. Mais recentemente, grandes modelos de linguagem (LLMs) têm sido empregados nessa tarefa (Gilardi et al., 2023; Tan et al., 2024), viabilizando inferências com poucos ou mesmo sem exemplos rotulados (*few-shot* e *zero-shot*). Contudo, tais métodos geralmente se restringem a um único *corpus* ou a um domínio específico e, em muitos casos, não avaliam sistematicamente a confiabilidade dos rótulos artificiais gerados. Além disso, o elevado custo computacional dos LLMs ainda representa um obstáculo relevante à adoção ampla dessas abordagens.

Diante desse cenário, este trabalho propõe um comitê constituído por três métodos complementares: (i) um método semi-supervisionado baseado em grafos heterogêneos, que explora relações estruturais entre textos e unidades linguísticas para a propagação de rótulos; (ii) uma estratégia fundamentada em aprendizagem *few-shot* com grandes modelos de linguagem; e (iii) um método baseado em *Retrieval-Augmented Generation* (RAG), que incorpora exemplos previamente anotados por meio de mecanismos de recuperação semântica. As decisões produzidas pelos três métodos são combinadas por meio de votação majoritária, constituindo o comitê proposto. Essa estratégia de agregação tem como objetivo aumentar a robustez do processo de anotação automática e reduzir a influência de vieses específicos de cada método, explorando a complementaridade entre abordagens fundamentadas em paradigmas distintos, cujos benefícios se manifestam de forma mais pronunciada em cenários de maior equilíbrio entre os métodos constituintes. A metodologia foi concebida para operar sob restrições de dados e em cenários de mudança de domínio, com o objetivo central de

¹<https://olhardigital.com.br/2023/01/19/pro/chatgpt-openai-explorou-trabalhadores-quenianos-para-identificar-conteudos-ilegais-e-criminosos/>

gerar rótulos automáticos confiáveis para a construção e ampliação de *corpora* anotados. A proposta parte de um único *corpus* curado, composto por 1.400 instâncias (Neto et al., 2024), utilizado como base de treinamento, e é avaliada em *corpora* amplamente empregados na literatura (Vargas et al., 2022; Fortuna et al., 2019; Leite et al., 2020), com volumes significativamente superiores. Esse desenho experimental permite investigar se um volume reduzido de dados anotados é suficiente para sustentar a geração automática de rótulos confiáveis em cenários de domínio cruzado.

Para avaliar, de forma sistemática, a robustez dessa estratégia sob diferentes condições de anotação, a metodologia é examinada em dois cenários experimentais complementares. No primeiro, o comitê é avaliado com base nos *corpora* em suas versões originais, preservando os rótulos atribuídos pelos anotadores humanos e, consequentemente, as divergências interanotador inerentes ao processo de anotação manual. No segundo cenário, os conjuntos de dados são filtrados para incluir apenas instâncias com concordância plena entre os anotadores, resultando em subconjuntos com menor divergência anotativa e, presumivelmente, maior consistência nos rótulos atribuídos. A comparação entre esses cenários permite analisar, de forma controlada, em que medida a consistência das anotações de referência influencia o desempenho, a concordância e a estabilidade dos métodos de anotação automática.

Nesse enquadramento, o presente trabalho organiza sua investigação em torno de quatro questões de pesquisa, que orientam tanto o desenho experimental quanto a análise dos resultados:

QP1 Entre os classificadores individuais que compõem o comitê, semi-supervisionado baseado em grafos heterogêneos, aprendizagem *few-shot* e *Retrieval Augmented Generation*, qual apresenta melhor desempenho na anotação automática de linguagem tóxica sob escassez de dados rotulados?

QP2 A combinação desses métodos por meio de votação majoritária, constituindo um comitê de classificadores, resulta em desempenho superior ao do melhor método individual nos diferentes *corpora* avaliados?

QP3 Como o desempenho dos classificadores automáticos, incluindo o comitê, varia ao considerar apenas instâncias com concordância plena entre os anotadores, mantendo as mesmas configurações experimentais, e o que isso indica sobre a consistência das anotações de referência?

QP4 Em que medida o grau de concordância entre os métodos individuais do comitê está associado aos ganhos obtidos pela estratégia de votação majoritária nos diferentes *corpora* e cenários experimentais avaliados?

A partir das investigações conduzidas em torno dessas questões, destacam-se como principais contribuições deste trabalho os seguintes aspectos:

- Proposta de um comitê de classificadores para a anotação automática de linguagem tóxica em português, baseada na integração de três paradigmas complementares, semi-supervisionado em grafos heterogêneos, aprendizagem *few-shot* e *Retrieval-Augmented Generation*, combinados por votação majoritária.
- Avaliação sistemática do comitê em cenário de domínio cruzado, com treinamento realizado em um único *corpus* curado de pequeno volume e aplicação a três *corpora* de avaliação com volumes substancialmente superiores e domínios distintos, permitindo verificar a capacidade de generalização da abordagem sob escassez de dados rotulados.
- Análise comparativa em dois cenários de anotação, versões originais dos *corpora* e subconjuntos filtrados por concordância plena entre anotadores, evidenciando o impacto da consistência das anotações de referência sobre o desempenho dos métodos automáticos avaliados.
- Avaliação de equidade no comportamento dos classificadores em relação a instâncias contendo termos identitários, complementando a análise de desempenho com indicadores de equidade fundamentados na literatura (Borkan et al., 2019; Dixon et al., 2018).
- Disponibilização da metodologia e dos resultados como contribuição para a expansão de recursos linguísticos anotados em português, língua ainda sub-representada no contexto de pesquisas em detecção e anotação automática de linguagem tóxica.

Este artigo está organizado da seguinte forma: a Seção 2 apresenta a fundamentação conceitual sobre toxicidade e fenômenos correlatos; a Seção 3 discute os trabalhos relacionados; a Seção 4 detalha os *corpora* utilizados; a Seção 5 descreve a metodologia proposta; a Seção 6 apresenta a configuração experimental, os resultados e a discussão; e, por fim, a Seção 7 discute as conclusões e direções para trabalhos futuros.

2. Toxicidade e Conceitos Correlatos

Na literatura, é possível encontrar uma variedade de definições para linguagem tóxica e seus correlatos, que dependem da perspectiva teórica sob a qual esse fenômeno é estudado. No contexto da sociolinguística, a linguagem tóxica é compreendida em sua dimensão performativa, na qual palavras ofensivas atuam como ações capazes de causar danos ou incitar à violência no mundo real (Butler, 2021). Nessa perspectiva, o significado da toxicidade não é determinado apenas pela estrutura léxica ou sintática do texto, mas também pela compreensão pragmática que articula os participantes do discurso e a história discursiva, configurando-a como um ato e não apenas como conteúdo. Em contraste, na literatura derivada do Processamento de Linguagem Natural, as definições acerca da linguagem tóxica tratam esse fenômeno de forma operacional, voltada à classificação automática de textos, organizando-o em manifestações específicas, como discurso de ódio, agressividade, linguagem ofensiva e *cyberbullying* (Fortuna & Nunes, 2018). Embora cada uma dessas categorias apresente particularidades relevantes para a tarefa de anotação, é frequente a sobreposição entre seus elementos característicos, o que pode introduzir ambiguidade no processo de anotação. O discurso de ódio refere-se especificamente a ataques baseados em características identitárias, como raça, gênero ou orientação sexual (Schmidt & Wiegand, 2017; Waseem & Hovy, 2016), enquanto a agressividade caracteriza-se pela intenção deliberada de agredir, sem necessariamente se vincular a marcadores identitários (Zampieri et al., 2019). A linguagem ofensiva, por sua vez, manifesta-se por meio de expressões obscenas ou depreciativas, sejam explícitas ou veladas (Fortuna & Nunes, 2018; Vargas et al., 2022), ao passo que o *cyberbullying* distingue-se pela frequência e persistência das agressões direcionadas a uma vítima específica (Mladenović et al., 2021). Essa sobreposição, somada à possibilidade de manifestações implícitas baseadas em ironia ou sarcasmo, pode comprometer a consistência dos *corpora* produzidos e, conseqüentemente, o desempenho dos modelos treinados a partir deles.

Nesse enquadramento teórico, o presente trabalho adota a definição proposta por Poletto et al. (2020), que organiza os diferentes fenômenos de linguagem abusiva em uma estrutura conceitual na qual a **toxicidade** é compreendida como o conceito mais abrangente, englobando manifestações como discurso de ódio, agressividade e linguagem ofensiva, conforme ilustrado na Figura 1. A adoção dessa estrutura

justifica-se em razão de sua compatibilidade com a maioria dos *corpora* disponíveis em português, cujos esquemas de anotação variam quanto à granularidade, e por permitir que o método proposto seja aplicado a diferentes domínios sem a necessidade de remapeamento conceitual, o que favorece sua aplicação em cenários de mudança de domínio. Reconhece-se, contudo, que essa escolha implica uma simplificação, uma vez que o tratamento de todas as manifestações ofensivas sob um único rótulo binário não captura subtipos específicos de toxicidade. Assim, ao longo deste artigo, toda manifestação de caráter ofensivo é tratada como linguagem tóxica, respeitadas as particularidades conceituais inerentes a cada categoria correlata.

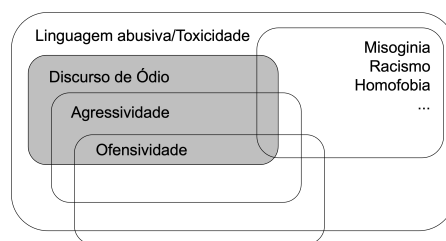


Figura 1: Estrutura conceitual de toxicidade e fenômenos de linguagem abusiva relacionados, adaptada de Poletto et al. (2020).

3. Trabalhos Relacionados

O combate ao conteúdo ofensivo em redes sociais tem motivado numerosos esforços acadêmicos voltados ao desenvolvimento de métodos automáticos para a detecção de linguagem tóxica (Fortuna & Nunes, 2018; Poletto et al., 2020). No contexto da língua portuguesa, diferentes abordagens têm sido investigadas, incluindo estratégias baseadas em léxicos (Pelle et al., 2018; Vargas et al., 2021), métodos tradicionais de Aprendizagem de Máquina (de Pelle & Moreira, 2017; Leite et al., 2020), modelos de *Deep Learning* (Soto et al., 2019) e classificadores fundamentados em arquiteturas *transformers* (Frediani et al., 2025). Apesar dos avanços alcançados, tais abordagens permanecem fortemente dependentes da disponibilidade de grandes volumes de dados anotados para treinamento, validação e avaliação. Mais recentemente, grandes modelos de linguagem também passaram a ser explorados na detecção de toxicidade, alcançando resultados expressivos (Oliveira et al., 2023, 2024; Assis et al., 2024). Contudo, essas abordagens ainda enfrentam limitações associadas aos elevados custos computacionais e aos desafios de escalabilidade, especialmente em cenários de aplicação em larga escala.

Embora os avanços na detecção automática de linguagem tóxica sejam expressivos, a dependência de grandes volumes de dados anotados permanece um dos principais gargalos da área. A anotação manual é extremamente onerosa, além de demandar muito tempo em sua realização. Além disso, está sujeita a vieses socio-culturais e à subjetividade interpretativa inerente à tarefa, o que pode comprometer a consistência e a qualidade dos dados e, consequentemente, afetar o desempenho dos métodos de detecção de toxicidade (Aroyo et al., 2019; Hettiachchi et al., 2023). Diante desse cenário, cresce o interesse por estratégias de anotação automática voltadas à criação e à ampliação de *corpora* linguísticos, com o objetivo de reduzir custos e acelerar o processo de anotação.

Entre essas estratégias, destacam-se abordagens baseadas em auto-treinamento, em que um classificador inicialmente treinado com um conjunto reduzido de dados anotados é utilizado para rotular automaticamente novas instâncias, posteriormente incorporadas de forma iterativa ao processo de treinamento. Nessa linha, Saifullah et al. (2024) propõem um método para anotação automática de discurso de ódio em comentários do YouTube em indonésio, combinando representações *TF-IDF* e *Word2Vec* em um esquema de meta-vetorização. De maneira semelhante, Alsafari & Sadaoui (2021) exploram o auto-treinamento no contexto da língua árabe, avaliando diferentes combinações de classificadores e adotando limiares elevados de confiança para incorporar instâncias automaticamente rotuladas. Apesar de demonstrarem potencial para a expansão de *corpora*, tais abordagens permanecem fortemente dependentes do desempenho inicial do modelo e dos critérios de seleção, o que pode resultar na propagação de erros e na amplificação de vieses ao longo do processo iterativo.

Avançando nessa direção, Santos et al. (2022) propõem uma abordagem que integra auto-treinamento, propagação de rótulos por similaridade semântica e um classificador baseado em *Generative Adversarial Networks* associado ao BERT, organizados em um esquema iterativo. A metodologia transfere conhecimento de *corpora* previamente anotados, originalmente coletados em outros domínios, para anotar automaticamente um *corpus* não rotulado em um cenário *cross-domain*. Os resultados indicam que, embora a estratégia seja eficaz para ampliar o volume de dados anotados automaticamente, o simples aumento do conjunto de treinamento não garante melhorias consistentes no desempenho, podendo, inclusive, introduzir ruído e vieses no

processo de anotação. Em outra vertente semi-supervisionada, Neto et al. (2024) investigam uma abordagem baseada em grafos heterogêneos para a anotação automática de linguagem tóxica. O estudo configura-se como uma investigação inicial, restrita a um único *corpus* e a uma configuração específica de *embeddings* estáticas, não explorando de forma sistemática cenários de mudança de domínio nem os desafios associados à detecção de formas implícitas de toxicidade, como ironia e sarcasmo, o que evidencia lacunas ainda existentes na modelagem da subjetividade linguística em português.

Estratégias de *pseudo-labeling* associadas a modelos supervisionados baseados em arquiteturas *transformers*, em especial as derivadas do BERT (Devlin et al., 2019), têm sido amplamente exploradas na literatura. Nessa direção, Suryawanshi et al. (2020) utilizam um conjunto reduzido de dados anotados manualmente para treinar diferentes classificadores, observando desempenho superior nas arquiteturas baseadas em BERT. O modelo selecionado é então empregado para rotular automaticamente um volume maior de dados não anotados, que passam a compor um conjunto expandido, utilizado em uma nova etapa de treinamento. De forma semelhante, Dirting et al. (2022) propõem uma abordagem também baseada em BERT para a classificação multi-rótulo da severidade do discurso de ódio, utilizando *pseudo-labeling* para anotar dados provenientes de múltiplas plataformas sociais. Embora esses estudos reportem melhorias consistentes em métricas tradicionais, como precisão, cobertura e *F-measure*, permanecem limitadas as análises acerca da confiabilidade dos rótulos automaticamente gerados e dos efeitos cumulativos da possível propagação de erros ao longo do processo iterativo de anotação.

Outros estudos investigam estratégias fundamentadas na similaridade semântica e no uso explícito de recursos linguísticos para a anotação automática de dados tóxicos. Phanomtip et al. (2021), por exemplo, exploram a combinação de dados rotulados e não rotulados na detecção de *cyberbullying* em tweets, empregando um classificador supervisionado baseado em SVM com representações *TF-IDF*, bem como métodos baseados em similaridade do cosseno. Os resultados indicam que a incorporação de instâncias automaticamente rotuladas por meio de um classificador supervisionado contribui para a melhoria do desempenho, enquanto abordagens fundamentadas exclusivamente na similaridade apresentam desempenho inferior. Em outra linha, Pelosi et al. (2017) propõem uma metodologia

para a anotação automática de linguagem ofensiva em italiano baseada em léxicos especializados e regras linguísticas, com ênfase na identificação de expressões tabu. Apesar dos resultados promissores, ambas as abordagens revelam forte dependência de recursos linguísticos específicos e enfrentam limitações na adaptação a novos domínios.

Após essas abordagens predominantemente automáticas, outras propostas buscam mitigar as limitações de escalabilidade e dependência de recursos específicos por meio de estratégias híbridas. Nesse sentido, abordagens semi-automáticas baseadas no paradigma *Human-in-the-Loop* têm sido propostas como forma de equilibrar custo e qualidade na anotação de dados textuais. Botella-Gil et al. (2024) integram técnicas de *Active Learning* a um processo de pré-anotação automática conduzido por modelos de aprendizado profundo, cujas previsões são posteriormente validadas por especialistas humanos. Aplicada à construção de um *corpus* em espanhol composto por mensagens violentas, a metodologia proporciona uma redução significativa no tempo de anotação manual, mantendo níveis competitivos de desempenho na detecção automática. Ainda assim, esse tipo de abordagem permanece dependente da disponibilidade contínua de anotadores humanos, o que pode comprometer sua escalabilidade em cenários de grande volume de dados.

Em síntese, embora a literatura registre avanços significativos no emprego de estratégias automáticas e semi-supervisionadas para a anotação de linguagem tóxica, persistem lacunas relevantes. Grande parte dos estudos são abordagens isoladas, avaliadas em cenários controlados e frequentemente restritas a um único *corpus* ou domínio específico, com análises limitadas quanto à capacidade de generalização entre contextos distintos. Ademais, as investigações que avaliam sistematicamente a confiabilidade dos rótulos gerados ainda são escassas, especialmente em tarefas caracterizadas por um elevado grau de subjetividade interpretativa, como a identificação de toxicidade implícita, de ironia e de sarcasmo. Permanecem igualmente pouco exploradas estratégias que articulem múltiplos paradigmas de anotação automática e que sejam avaliadas em diferentes níveis de ruído de anotação e em cenários de domínio cruzado. No contexto da língua portuguesa, tais limitações tornam-se ainda mais evidentes, em razão da escassez de recursos linguísticos anotados e da predominância de modelos supervisionados cuja capacidade de generalização é insuficientemente avaliada.

4. Corpora

Para o desenvolvimento e a avaliação da metodologia proposta, foram utilizados quatro *corpora* linguísticos de linguagem tóxica em língua portuguesa: Toxic-BR (Neto et al., 2024), ToLD-BR (Leite et al., 2020), HateBR (Vargas et al., 2022) e HLPHSD (Fortuna et al., 2019). A seleção desses conjuntos de dados buscou contemplar diferentes contextos de coleta, esquemas de anotação e plataformas de origem, permitindo uma avaliação mais abrangente da robustez e da capacidade de generalização da abordagem proposta em cenários de mudança de domínio.

O *corpus* Toxic-BR é composto por 1.400 textos em língua portuguesa extraídos da rede social *X* (anteriormente *Twitter*), coletados durante o segundo turno das eleições presidenciais brasileiras de 2022, período marcado por intensa polarização política e por frequente troca de ofensas entre apoiadores dos candidatos (Neto et al., 2024). Inicialmente, os dados foram automaticamente anotados de forma binária (*tóxico* e *não tóxico*) por meio da *Perspective API* (Lees et al., 2022) e do modelo GPT-3 (Brown et al., 2020); posteriormente, essas anotações passaram por um processo de curadoria manual conduzido por especialistas humanos, que avaliaram a qualidade das classificações automáticas e corrigiram inconsistências, visando aumentar a consistência e a confiabilidade dos rótulos finais.

Além do Toxic-BR, este trabalho utiliza o ToLD-BR, um *corpus* composto por 21.000 *tweets* manualmente anotados em sete categorias, incluindo racismo, linguagem obscena, insulto, xenofobia, LGBTQ+fobia, misoginia e conteúdos não tóxicos (Leite et al., 2020). O processo de anotação envolveu 48 colaboradores, organizados em grupos de três avaliadores responsáveis por aproximadamente 1.500 textos cada, o que permitiu múltiplas avaliações por instância. Adicionalmente, os autores disponibilizam uma versão binária do conjunto, na qual 9.255 mensagens originalmente associadas às categorias de toxicidade foram agrupadas na classe tóxica, enquanto as demais foram rotuladas como não tóxicas.

De maneira complementar, o *corpus* HateBR é constituído por 7.000 textos em português, coletados a partir de publicações em perfis políticos da rede social *Instagram* (Vargas et al., 2022). As anotações contemplam diferentes esquemas, incluindo classificação binária (ofensivo vs. não ofensivo), níveis graduais de ofensividade e categorias específicas de discurso de ódio, como sexismo, homofobia, intolerância religiosa e racismo. Um diferencial desse conjunto reside no

fato de ter sido integralmente anotado por três especialistas com elevada formação acadêmica, o que contribui para maior consistência e qualidade nas rotulações.

Por fim, o *corpus* HLPHSD também reúne textos oriundos da antiga rede social *Twitter*, totalizando 3.882 *tweets* neutros e 1.788 *tweets* rotulados como discurso de ódio (Fortuna et al., 2019). Quando identificado conteúdo ofensivo, são atribuídos rótulos hierárquicos adicionais correspondentes a categorias específicas, como sexismo, homofobia e racismo. Para a tarefa de classificação binária, cada *tweet* foi avaliado por três anotadores, o que permitiu analisar o grau de concordância entre os avaliadores humanos.

A Tabela 1 sintetiza a distribuição de instâncias tóxicas e não tóxicas em cada *corpus*, bem como a média de *tokens* por texto. De modo geral, observa-se que a classe não tóxica é a mais frequente em todos os conjuntos, o que reflete tanto a distribuição natural de conteúdos em ambientes digitais quanto os desafios inerentes à coleta e à anotação de mensagens explicitamente ofensivas. Além disso, nota-se que os *corpora* oriundos da rede social *X* apresentam, em média, textos mais longos do que os do HateBR, possivelmente em razão de diferenças no estilo discursivo e na dinâmica de interação entre as plataformas.

<i>Corpus</i>	Tóxicos	Não Tóxicos	Origem	#Tokens
Toxic-BR	615	785	<i>X</i>	49.47
ToLD-BR	9.255	11.745	<i>X</i>	30.22
HateBR	3.500	3.500	<i>Instagram</i>	24.31
HLPHSD	1.788	3.882	<i>X</i>	37.43

Tabela 1: Distribuição das instâncias por classe, origem dos textos, bem como a média do número de *tokens* por instância nos *corpora* em suas versões originais.

Além das características de distribuição apresentadas, os *corpora* analisados apresentam níveis distintos de concordância interanotador, conforme relatado nos respectivos trabalhos. Tais variações, no entanto, não são interpretadas neste estudo como evidência direta de maior ou menor ambiguidade entre os conjuntos de dados, uma vez que podem decorrer de diferenças nos protocolos de anotação, nas diretrizes adotadas, no perfil dos anotadores ou nos procedimentos de validação dos rótulos. Considerando, contudo, que a identificação de linguagem tóxica constitui uma tarefa intrinsecamente subjetiva, instâncias com concordância total entre os anotadores são, neste trabalho, tratadas como casos de menor divergência interanotador no contexto específico de cada *corpus*. Nesse sentido, realizou-se uma

etapa adicional de filtragem, mantendo-se apenas os textos para os quais houve consenso absoluto quanto ao rótulo atribuído. A Tabela 2 apresenta a distribuição das instâncias após esse procedimento, evidenciando, como efeito direto, um aumento do desbalanceamento entre as classes e uma redução das médias de *tokens* dos *corpora* filtrados. O uso específico de cada conjunto na avaliação experimental é detalhado na Seção 6.

<i>Corpus</i>	Tóxicos	Não Tóxicos	#Tokens
ToLD-BR	1.453	11.571	21.12
HateBR	2.382	3.298	18.17
HLPHSD	507	1.957	19.75

Tabela 2: Distribuição de instâncias tóxicas e não tóxicas, bem como a média do número de *tokens* por texto, nos *corpora* após a filtragem por concordância total.

5. Metodologia

Esta Seção descreve a metodologia desenvolvida para a anotação automática de linguagem tóxica. O comitê de classificadores integra três abordagens complementares que compartilham o mesmo objetivo de atribuição automática de rótulos de toxicidade, mas exploram princípios distintos de modelagem e inferência: um método semi-supervisionado responsável pela propagação de rótulos em grafos heterogêneos; um classificador baseado em LLM que realiza inferência orientada por exemplos; e outro classificador que conduz a inferência enriquecida por mecanismos de recuperação semântica. Essa integração tem como objetivo aumentar a robustez do processo de anotação automática e reduzir a influência de vieses específicos de cada método, explorando a complementaridade entre abordagens com princípios distintos de modelagem e inferência.

A Figura 2 apresenta a visão geral da metodologia. O processo inicia-se com textos brutos, submetidos a pré-processamento e normalização. O texto resultante é então encaminhado aos três métodos do comitê, que realizam, de forma independente, a classificação binária das instâncias como *tóxicas* ou *não tóxicas*. As predições são agregadas por meio de votação majoritária, o que define o rótulo final. Os textos automaticamente anotados podem ser utilizados na construção ou na ampliação de *corpora* de linguagem tóxica. Embora tradicionalmente associadas à detecção, as abordagens são empregadas sob a perspectiva da anotação automática.

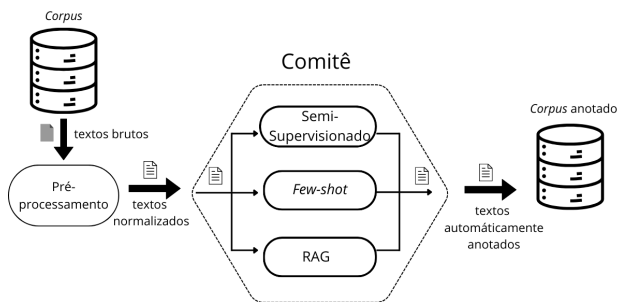


Figura 2: Visão geral da metodologia com um comitê de classificadores para a anotação automática de textos tóxicos.

Nas Subseções seguintes, detalham-se a etapa de pré-processamento (Subseção 5.1), os métodos que compõem o comitê (Subseções 5.2, 5.3 e 5.4) e a estratégia de combinação das predições (Subseção 5.5).

5.1. Pré-processamento

A etapa de pré-processamento consiste na normalização do conteúdo textual utilizado pelos métodos do comitê. Para isso, foram removidas menções a *URLs*, indicadores de *retweets* (RT), referências a usuários e *emojis*, e foram descartados os textos compostos exclusivamente por *emojis*. Além disso, considerando a predominância de linguagem coloquial em ambientes digitais, caracterizada por abreviações e variações ortográficas frequentes, aplicou-se a ferramenta Enelvo (Costa Bertaglia & Volpe Nunes, 2016) para a normalização textual. Esse procedimento contribui para tornar os dados mais homogêneos e, conseqüentemente, mais adequados ao processamento pelos classificadores do comitê.

5.2. Método Semi-Supervisionado

O método semi-supervisionado adotado neste trabalho foi inspirado na proposta de (Neto et al., 2024), que modela textos tóxicos por meio de grafos heterogêneos e emprega um algoritmo de propagação de rótulos para a anotação automática de linguagem tóxica. Para sua integração ao comitê, a abordagem foi estendida com a incorporação de *word embeddings* contextuais e com a aplicação conjunta do algoritmo original de propagação de rótulos e de uma estratégia alternativa, permitindo examinar diferentes dinâmicas de propagação em grafos, aspectos não explorados no estudo original.

O método semi-supervisionado está estruturado em três etapas principais: (i) modelagem do grafo, (ii) regularização e (iii) classificação, as quais são descritas a seguir.

5.2.1. Modelagem do grafo

Grafos constituem estruturas amplamente empregadas na representação de dados e têm recebido atenção crescente na última década, sendo aplicados com êxito em diversas tarefas, como a modelagem de tópicos e a desambiguação de nomes, frequentemente com resultados promissores (King et al., 2014). Em particular, grafos heterogêneos possibilitam a integração de informações estruturais, expressas por arestas que conectam nós de diferentes tipos com conteúdos não estruturados associados a cada entidade, resultando em modelos mais expressivos e semanticamente enriquecidos (Zhang et al., 2019). Sua principal contribuição reside na explicitação das relações entre entidades distintas, o que favorece uma representação mais informativa e contextualizada em múltiplos domínios de aplicação.

Com base nesses princípios, a primeira etapa do método semi-supervisionado consiste na modelagem dos textos em um grafo heterogêneo, no qual cada instância é representada por nós correspondentes à sentença textual, aos *tokens* que a compõem e a um nó adicional denominado Toxicidade. Esse nó não representa o rótulo final da instância, mas atua como uma característica auxiliar incorporada à estrutura do grafo, contribuindo exclusivamente para a modelagem relacional durante o processo de regularização, sem ser utilizado como variável supervisionada nem como rótulo de saída do método. O valor associado ao nó Toxicidade corresponde a um índice de ofensividade estimado com base nos termos potencialmente ofensivos presentes no texto, identificados por meio do léxico MOL (Vargas et al., 2021). O grau de toxicidade de cada termo é definido com o auxílio da *Perspective API* (Lees et al., 2022), e o valor final da toxicidade do texto é calculado como a soma dos índices atribuídos aos termos ofensivos identificados.

No que se refere às arestas do grafo, estas são não direcionadas e ponderadas, estabelecendo conexões exclusivamente entre nós do tipo sentença e nós do tipo *token*, sem ligações diretas entre sentenças nem entre *tokens*. As conexões entre nós de sentença e o nó de Toxicidade, por sua vez, não recebem pesos e são utilizadas apenas para representar relações estruturais na modelagem do grafo. O peso de cada aresta que conecta um nó *sentença* a um nó *token* é calculado a partir da média dos valores do vetor de *embedding* associado ao termo correspondente, de modo que cada aresta passe a refletir um valor escalar que representa a contribuição semântica do *token* no contexto da sentença. Adicionalmente, nós de **sentença** compartilham nós de

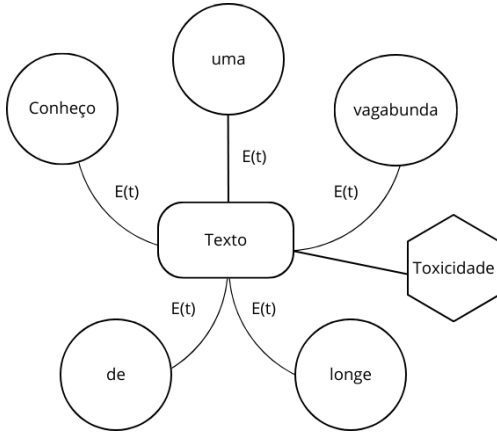


Figura 3: Exemplo da estrutura do grafo modelada para a sentença “Conheço uma vagabunda de longe!”.

token sempre que o mesmo termo ocorre em diferentes textos, permitindo que sentenças distintas se conectem indiretamente por meio desses *tokens* compartilhados e favorecendo a propagação relacional de informação ao longo da estrutura do grafo. A Figura 3 ilustra a configuração completa do grafo heterogêneo construído a partir de um exemplo de texto tóxico, evidenciando os nós de sentença, *tokens* e Toxicidade, bem como as arestas ponderadas que conectam sentenças e *tokens*. Diferentemente da proposta original, a atribuição dos pesos das arestas foi realizada com base em *embeddings* contextuais extraídos do modelo BERTimbau (Souza et al., 2020), em sua versão pré-treinada, aplicado diretamente aos dados brutos, sem ajuste fino ou engenharia de *prompts*. Em contraste com *embeddings* distribucionais estáticos, representações contextuais, como as geradas pelo BERTimbau, capturam variações semânticas dependentes do contexto, produzindo vetores distintos para cada ocorrência de um mesmo *token* (Smith, 2020).

5.2.2. Regularização

A etapa de regularização é responsável pela extração de características a partir da estrutura do grafo, configurando-se como um processo de classificação transdutiva ou semi-supervisionada. Seu objetivo consiste em determinar um conjunto de rótulos que satisfaça simultaneamente duas propriedades: (i) consistência com os dados previamente rotulados e (ii) coerência com a topologia do grafo, sob a premissa de que nós vizinhos tendem a compartilhar rótulos semelhantes (Rossi, 2015).

Na proposta original de Neto et al. (2024), empregou-se exclusivamente o algoritmo *Gaussian Fields and Harmonic Functions* (GFHF) apresentado por Zhu et al. (2003). No presente trabalho, além do GFHF, incorpora-se também o algoritmo *Local and Global Consistency* (LGC) (Zhou et al., 2004), ampliando a análise para diferentes estratégias de propagação de rótulos. Ambos os métodos propagam rótulos no grafo heterogêneo com base na similaridade entre nós conectados; contudo, diferem no tratamento de nós previamente rotulados. Enquanto o GFHF mantém os rótulos iniciais fixos ao longo de todo o processo, o LGC permite sua atualização iterativa, conferindo maior flexibilidade ao mecanismo de propagação.

Para a execução dos algoritmos de regularização, é necessário definir um conjunto inicial de nós rotulados, que serve como ponto de partida para a classificação transdutiva. Por exemplo, ao selecionar 20% dos nós como pré-rotulados, pode-se assegurar que 10% das instâncias de cada classe sejam escolhidas aleatoriamente. A inferência dos rótulos dos nós não rotulados é realizada com base na similaridade com seus vizinhos mais próximos, ponderada pelos pesos das arestas do grafo. Como resultado, os algoritmos de regularização produzem, para cada nó, uma representação vetorial induzida pela estrutura relacional do grafo, conforme ilustrado na Tabela 3. Nessa tabela, o campo **ID** identifica o nó correspondente à instância textual; os campos **Valor 1** e **Valor 2** representam suas coordenadas; e o campo **Rótulo** indica a classe inferida pelo regularizador, sendo **1** para textos tóxicos e **0** para textos não tóxicos.

ID	Valor 1	Valor 2	Rótulo
440	0.03221	0.01447	1
702	0.05716	0.02365	1
6437	0.0029	0.00011	0
6464	-30.09728	-1.11603	0

Tabela 3: Exemplo de saída do algoritmo de regularização LGC.

5.2.3. Classificação

Na etapa de classificação, as coordenadas produzidas pelo processo de regularização (ver Tabela 3) são utilizadas como atributos de entrada para algoritmos de Aprendizagem de Máquina Supervisionada.

5.3. Método baseado em aprendizagem *few-shot*

A abordagem de aprendizagem *few-shot* orienta a inferência do LLM por meio da inclusão explícita de exemplos previamente anotados no *prompt*. Para cada conjunto de textos a ser classificado, são selecionadas quatro instâncias do *corpus* de treinamento, duas de cada classe, incorporadas como referências para a decisão sobre novas entradas e mantidas fixas ao longo de todo o processo, assegurando uniformidade na instrução fornecida ao modelo. A escolha dessa configuração preserva o comprimento do *prompt* dentro de limites compatíveis com modelos de menor porte (Brown et al., 2020), além de que a distribuição balanceada entre categorias contribui para reduzir o viés de instrução (da Silva Oliveira et al., 2024; Szcz et al., 2025). Os exemplos selecionados, juntamente com o texto-alvo, estruturam o *prompt* conforme o formato apresentado no Exemplo 5.3, que é então submetido ao modelo, que realiza a inferência e retorna a classificação binária da instância como *tóxico* ou *não tóxico*.

Exemplo (Prompt do Comitê). Você é um especialista em detecção de toxicidade em textos. A linguagem tóxica é caracterizada por qualquer forma de expressão fortemente indelicada, rude ou ofensiva, incluindo palavrões ou declarações que demonstrem desrespeito a indivíduos ou grupos.

Sua tarefa é analisar o texto fornecido e classificá-lo como "tóxico" ou "não_tóxico", considerando apenas o conteúdo textual apresentado e os exemplos fornecidos.

Saída:

- Responda exclusivamente em JSON.
- Não produza texto adicional.

Formato da saída:

```
{"label": <tóxico|nao_tóxico>"}
```

Exemplos: {exemplos}

Texto a classificar: "{texto}"

5.4. Método baseado em *Retrieval-Augmented Generation*

O método baseado em *Retrieval-Augmented Generation* (RAG) estende a abordagem *few-shot* ao substituir a seleção fixa de exemplos por um mecanismo automático de recuperação semântica (Lewis et al., 2020). Nessa estratégia, o texto-alvo é inicialmente convertido em uma representação vetorial e utilizado como consulta para recuperar os k exemplos mais semelhantes a partir de uma base indexada construída com o con-

junto de treinamento. A indexação e a busca vetorial são realizadas por meio da biblioteca FAISS² em conjunto com o modelo de *embeddings* *paraphrase-multilingual-MiniLM-L12-v2*³, selecionado por seu suporte multilíngue, incluindo o português, e pelo equilíbrio entre a qualidade de representação semântica e a eficiência computacional. Em seguida, os exemplos recuperados, juntamente com seus respectivos rótulos, são incorporados à construção do *prompt*, preservando o mesmo formato adotado no método *few-shot*. O *prompt* resultante é posteriormente encaminhado ao LLM para a etapa de inferência, na qual o rótulo de toxicidade é atribuído ao texto-alvo. O valor de $k = 4$ foi mantido para permitir uma comparação controlada entre as estratégias *few-shot* e RAG.

5.5. Estratégia de Combinação do Comitê

Após a execução independente dos três métodos descritos nas subseções anteriores, suas predições são agregadas por meio de votação majoritária. A predição final é definida pela Equação 1:

$$\hat{y} = \arg \max_j \sum_{t=1}^T d_{t,j} \quad (1)$$

em que $d_{t,j} = 1$ se o classificador t prediz a classe j , e $d_{t,j} = 0$ caso contrário. Essa formulação explícita que a classe selecionada é a que acumula o maior número de votos entre os classificadores do comitê. Essa estratégia de aprendizagem baseada em comitês é amplamente reconhecida por sua capacidade de aumentar a robustez e a estabilidade das predições ao combinar classificadores com comportamentos distintos (Zhou, 2012).

6. Experimentos e Resultados

Esta Seção apresenta as configurações experimentais adotadas para a avaliação do comitê de classificadores e para a análise dos resultados obtidos, com ênfase no desempenho da classe tóxica, por se tratar do caso mais crítico no contexto da detecção e da anotação automática de linguagem abusiva. Para essa finalidade, o desempenho do comitê e de seus métodos constituintes é avaliado por meio da métrica *F-measure* e do coeficiente kappa de Cohen. A *F-measure* foi escolhida por corresponder à média harmônica

²<https://python.langchain.com/docs/integrations/vectorstores/faiss/>

³<https://huggingface.co/sentence-transformers/paraphrase-multilingual-MiniLM-L12-v2>

entre precisão e cobertura, permitindo uma avaliação equilibrada entre a proporção de instâncias corretamente classificadas e a capacidade do modelo de recuperar casos relevantes. Complementarmente, o coeficiente Kappa de Cohen é empregado para mensurar o grau de concordância entre dois avaliadores, considerando tanto a concordância observada quanto a esperada ao acaso, o que resulta em uma estimativa mais realista da confiabilidade dos rótulos atribuídos; sua interpretação permite caracterizar níveis de concordância que variam de fraca a quase perfeita (Landis & Koch, 1977). Neste trabalho, o Kappa é calculado com base na concordância entre os rótulos originais dos *corpora* e os atribuídos automaticamente pelos métodos do comitê.

6.1. Configuração Experimental

Os métodos que compõem o comitê foram avaliados de forma independente, adotando-se o *corpus* Toxic-BR como base de treinamento e os *corpora* ToLD-BR, HateBR e HLPHSD exclusivamente como conjuntos de teste. Ressalta-se que o Toxic-BR reúne 1.400 instâncias (Neto et al., 2024), enquanto os conjuntos de avaliação variam entre 5.670 e 21.000 textos (Fortuna et al., 2019; Vargas et al., 2022; Leite et al., 2020). Essa assimetria configura um regime de treinamento com baixo volume de dados em relação aos cenários de avaliação em larga escala, permitindo examinar a capacidade de generalização do comitê a partir de um conjunto reduzido de exemplos previamente curados. Consequentemente, o desenho experimental permite analisar o comportamento dos métodos em diferentes domínios, plataformas de coleta e esquemas de anotação, sem necessidade de ajustes específicos para cada *corpus*. Ademais, como linhas de base comuns a todas as abordagens, considerou-se uma configuração *zero-shot*, com o propósito de mensurar o patamar de desempenho obtido exclusivamente por meio de *prompting*, sem o fornecimento de exemplos rotulados, e um classificador BERTimbau (Souza et al., 2020) submetido a ajuste fino no *corpus* Toxic-BR, permitindo situar os resultados do comitê em relação a um método supervisionado baseado em *transformers* (Wolf et al., 2020).

Com base nesse desenho experimental, definiram-se dois cenários de avaliação. No primeiro, os métodos foram aplicados às versões originais dos *corpora*, conforme apresentado na Tabela 1. No segundo, utilizaram-se apenas os subconjuntos compostos por instâncias com concordância total entre os anotadores humanos, visando reduzir o impacto do ruído de anotação

e investigar seus efeitos sobre o desempenho dos métodos (ver Tabela 2). Cabe ressaltar que esses dois cenários desempenham papéis complementares. As QP1 e QP2 são respondidas com base no cenário original. A QP3 é examinada exclusivamente no cenário filtrado, por tratar do efeito da concordância plena entre anotadores sobre o desempenho dos métodos. A QP4 é analisada em ambos os cenários, o que permite verificar como o grau de concordância entre os métodos do comitê se comporta entre os dois regimes de definição do rótulo.

No método semi-supervisionado baseado em grafos heterogêneos, a construção do grafo integra instâncias do conjunto de treinamento (Toxic-BR) como exemplos iniciais para a aplicação dos algoritmos de regularização, viabilizando a propagação de rótulos nos *corpora* de teste (ToLD-BR, HateBR e HLPHSD) e permitindo a exploração de relações estruturais entre textos provenientes de diferentes domínios. Conforme discutido em (Neto et al., 2024), a definição de um volume adequado de rótulos iniciais é fundamental para assegurar estatísticas mais estáveis na identificação de linguagem tóxica e maior alinhamento com os conjuntos de avaliação. Assim, foram investigadas proporções de 10%, 20% e 30% de instâncias pré-rotuladas. As representações vetoriais obtidas ao final da etapa de regularização foram empregadas como atributos de entrada para três classificadores supervisionados, implementados com a biblioteca *Scikit-Learn* (Pedregosa et al., 2011): um *Multi-Layer Perceptron* (MLP)⁴, uma *Support Vector Machine* (SVM) treinada por meio do *SGD-Classifer*⁵ e um modelo de *Gradient Boosting* (GB), implementado com o *HistGradientBoostingClassifier*⁶. Como a seleção das instâncias pré-rotuladas ocorre aleatoriamente, diferentes execuções podem gerar variações na estrutura do grafo e, consequentemente, nos resultados obtidos. Por essa razão, cada configuração experimental foi executada 10 vezes, considerando ambos os algoritmos de regularização, GFHF e LGC.

Já para os métodos baseados em grandes modelos de linguagem, incluindo o método de linha de base *zero-shot*, a estratégia *few-shot* e a abordagem fundamentada em *Retrieval-Augmented Generation*, optou-se por modelos de porte reduzido, considerando as limitações computaci-

⁴Configurado com 100 unidades na camada oculta, função de ativação ReLU, otimizador Adam, taxa de aprendizado de 0,001 e máximo de 300 iterações.

⁵Empregada com núcleo linear.

⁶Utilizado com parâmetros padrão da biblioteca, sem ajuste adicional.

onais disponíveis e o propósito de viabilizar a anotação automática em cenários de custo restrito. Nesse contexto, foram empregados os modelos multilíngues Qwen2.5-7B (Yang et al., 2025) e Granite3.3-8B (IBM Research, 2025), utilizados sem ajuste fino e explorados exclusivamente por meio de engenharia de *prompt*. Os experimentos com LLMs foram conduzidos localmente em uma máquina equipada com GPU *NVIDIA GeForce RTX 3060* (12 GB), CPU *Intel(R) Core(TM) i5-10400F* (2,90 GHz) e 128 GB de memória RAM. Para tanto, a implementação foi realizada com a biblioteca *Ollama*⁷, que permite a execução eficiente de modelos de linguagem em ambiente local.

Como linha de base adicional baseada em aprendizado supervisionado, avaliou-se o modelo BERTimbau *Base* (Souza et al., 2020) submetido a ajuste fino no *corpus* Toxic-BR. O modelo foi configurado para classificação binária, com uma camada linear aplicada ao vetor do token [CLS], e treinado por 3 épocas, com tamanho de batch de 16, utilizando a biblioteca *transformers* da *Hugging Face* (Wolf et al., 2020)⁸. Manteve-se o mesmo cenário de domínio cruzado adotado para os demais métodos, com treinamento exclusivo no Toxic-BR e a avaliação nos *corpora* ToLD-BR, HateBR e HLPDSD, sem qualquer ajuste fino adicional nos conjuntos de avaliação.

6.2. Resultados Globais dos Classificadores

Esta Subseção apresenta uma visão geral do desempenho dos métodos avaliados, fornecendo uma base comparativa para a discussão detalhada nas Subseções seguintes. A Tabela 4 reúne os melhores resultados obtidos por cada classificador individual nos três *corpora* de avaliação, em termos de *F-measure* para a classe tóxica e do coeficiente Kappa de Cohen, considerando os *corpora* em suas versões originais.

De modo geral, os resultados apresentados na Tabela 4 indicam que os classificadores baseados em LLMs apresentam desempenho superior ao do classificador semi-supervisionado e ao classificador baseado em BERTimbau com ajuste fino na maioria dos cenários avaliados, sendo essa diferença mais pronunciada nos *corpora* HateBR e HLPDSD, enquanto, no ToLD-BR, os resultados das abordagens são mais próximos entre si. Entre as estratégias baseadas em LLMs, o classi-

Corpus	Método	Configuração	F	κ
ToLD-BR	BERT	BERTimbau <i>ajuste-fino</i>	0,72	0,45
	<i>zero-shot</i>	Granite3.3-8B	0,73	0,47
	semi	LGC_10%_SVM	0,72	0,46
	<i>few-shot</i>	Granite3.3-8B	0,74	0,48
	RAG	Qwen2.5-7B	0,74	0,45
HateBR	BERT	BERTimbau <i>ajuste-fino</i>	0,82	0,66
	<i>zero-shot</i>	Qwen2.5-7B	0,86	0,71
	semi	LGC_30%_MLP	0,37	0,22
	<i>few-shot</i>	Qwen2.5-7B	0,87	0,73
	RAG	Qwen2.5-7B	0,85	0,70
HLPDSD	BERT	BERTimbau <i>ajuste-fino</i>	0,56	0,37
	<i>zero-shot</i>	Qwen2.5-7B	0,62	0,38
	semi	LGC_10%_SVM	0,41	0,06
	<i>few-shot</i>	Qwen2.5-7B	0,63	0,41
	RAG	Qwen2.5-7B	0,61	0,39

Tabela 4: Melhores resultados obtidos por cada método nos *corpora* avaliados em suas versões originais, em termos de *F-measure* (F) e coeficiente Kappa de Cohen (κ). A configuração adotada em cada caso é indicada na coluna **Configuração**. Valores em negrito indicam o melhor resultado global em cada *corpus*.

ficador *few-shot* obteve os melhores resultados individuais em dois dos três *corpora* avaliados (HateBR e HLPDSD), além de apresentar desempenho equivalente ao do RAG no ToLD-BR e de apresentar ganhos em relação aos classificadores de linha de base, tanto *zero-shot* quanto BERTimbau com ajuste fino, nos três *corpora*.

O classificador RAG, por sua vez, apresenta desempenho comparável ao *zero-shot* em parte dos cenários e superior ao BERTimbau, mantendo comportamento competitivo entre as abordagens baseadas em LLM. Esses resultados são observados em configurações que utilizam exemplos rotulados no *prompt*, seja por seleção fixa (*few-shot*) ou por recuperação semântica (RAG), nos cenários de melhor desempenho. No caso do classificador semi-supervisionado, observa-se desempenho competitivo apenas no *corpus* ToLD-BR, com resultados próximos aos das abordagens baseadas em LLM, enquanto nos demais *corpora* os valores de *F-measure* e Kappa são significativamente inferiores, evidenciando maior sensibilidade do método à variação de domínio.

Quanto à **QP1**, o classificador *few-shot* apresentou o melhor desempenho individual na anotação automática de linguagem tóxica nos *corpora* avaliados, com o resultado mais expressivo observado no *corpus* HateBR no cenário original.

⁷<https://ollama.com/>

⁸Foram utilizados a taxa de aprendizado de (5×10^{-5}) , aquecimento inicial de 500 passos e decaimento de peso de 0,01.

6.3. Análise da Anotação Automática por Classificador

As subseções a seguir detalham o comportamento de cada classificador constituinte do comitê, com ênfase em aspectos que não são diretamente capturados pela Tabela 4, como padrões de variação entre as configurações, sensibilidade à mudança de domínio e diferenças entre os modelos de linguagem avaliados. Os resultados das melhores configurações do método semi-supervisionado e das abordagens baseadas em LLMs, detalhados pelos modelos Qwen2.5-7B e Granite3.3-8B para cada método e *corpora*, são apresentados nos Apêndices A e B (Tabelas 11 e 12).

6.3.1. Linhas de Base: Zero-Shot e BERTimbau

Dois métodos foram adotados como linhas de base para comparação com os classificadores do comitê. O primeiro, o método *zero-shot*, emprega os mesmos modelos de LLMs e a mesma estrutura de *prompt* adotados nas abordagens *few-shot* e RAG, distinguindo-se apenas pela ausência de exemplos rotulados. O segundo, o classificador BERTimbau com ajuste fino no Toxic-BR, representa uma abordagem supervisionada baseada em *transformers*, permitindo situar os resultados do comitê em relação a um paradigma amplamente consolidado na literatura para detecção de toxicidade (Leite et al., 2020; Oliveira et al., 2024; Salles et al., 2025).

No que se refere ao *zero-shot*, os modelos Qwen2.5-7B e Granite3.3-8B apresentam desempenho semelhante entre si nos três *corpora* avaliados. Já o BERTimbau permanece sistematicamente abaixo das demais abordagens avaliadas, incluindo o *zero-shot*, com diferença mais pronunciada no HLPHSD, conforme indicado na Tabela 4, comportamento esperado em um cenário de domínio cruzado no qual o ajuste fino foi conduzido em um único *corpus* de domínio específico.

6.3.2. Classificador Semi-Supervisionado

Para o método semi-supervisionado, as melhores configurações por *corpus* foram selecionadas com base no desempenho obtido em cada conjunto, a partir das combinações descritas na Subseção 6.1. O algoritmo LGC com classificador SVM e 10% de dados pré-rotulados apresentou os melhores resultados nos *corpora* ToLD-BR e HLPHSD, enquanto a combinação entre LGC, classificador MLP e 30% de dados pré-rotulados obteve o melhor desempenho no *corpus* HateBR. Cumpre destacar que todas as melhores configurações par-

tiram do algoritmo LGC, evidenciando que, no contexto de domínio cruzado e de identificação de toxicidade aqui investigado, a possibilidade de atualização iterativa dos rótulos pré-anotados se mostrou mais eficaz do que a manutenção fixa adotada pelo GFHF, ainda que os ganhos do LGC nos *corpora* HateBR e HLPHSD permaneçam aquém de patamares satisfatórios. Essas configurações foram adotadas como referência para os experimentos subsequentes.

Considerando essas configurações, observa-se que o classificador semi-supervisionado apresenta desempenho mais elevado no *corpus* ToLD-BR, com *F-measure* de 0,72 e níveis de concordância moderados. Nos *corpora* HateBR e HLPHSD, por sua vez, os valores de desempenho são inferiores, com redução tanto na *F-measure* quanto nos níveis de concordância. Em comparação aos demais classificadores avaliados, o semi-supervisionado mostra-se competitivo no ToLD-BR, com resultados próximos aos do BERTimbau, do *zero-shot* e das abordagens baseadas em LLMs com exemplos rotulados. Nos *corpora* HateBR e HLPHSD, contudo, fica sistematicamente abaixo de todos eles. Esse comportamento ocorre em cenários de mudança de domínio entre o conjunto de treinamento e os dados avaliados, condição associada à redução da concordância com as anotações humanas nesses *corpora*. Enquanto no ToLD-BR os valores de Kappa indicam concordância moderada, nos demais *corpora* o nível de concordância varia de fraco a razoável, evidenciando limitações do método na reprodução dos padrões de rotulagem em cenários de domínio cruzado.

6.3.3. Classificadores baseados em LLMs: *few-shot* e RAG

As abordagens *few-shot* e RAG apresentam comportamentos próximos entre si nos três *corpora* avaliados quando se considera o desempenho agregado dos dois métodos, com ganhos discretos em relação à linha de base *zero-shot* e desempenho consistentemente superior ao do BERTimbau com ajuste fino, em particular nos *corpora* HateBR e HLPHSD. No ToLD-BR, ambos os métodos apresentam desempenho equilibrado, com concordância moderada em relação aos rótulos humanos. No HateBR, observa-se o melhor desempenho global do estudo, alcançado pelo *few-shot* com o modelo Qwen2.5-7B, com concordância substancial. No HLPHSD, o desempenho de ambas as abordagens é mais modesto, sugerindo que a incorporação de exemplos no *prompt* é menos eficaz nesse domínio.

No que se refere às nuances entre os modelos, o Qwen2.5-7B apresenta melhor desempenho na configuração *few-shot* do que na configuração RAG nos três *corpora* avaliados. Em contrapartida, com o Granite3.3-8B, observa-se o padrão inverso nos *corpora* HateBR e HLPHSD, em que a configuração RAG supera o *few-shot* com o mesmo modelo, ainda que, no ToLD-BR, o *few-shot* permaneça superior. Esse resultado evidencia que as diferenças entre as duas abordagens não são uniformes entre os modelos e os *corpora* avaliados. De modo geral, as diferenças entre *few-shot* e RAG permanecem marginais, indicando que ambas as estratégias produzem inferências semelhantes quando aplicadas ao mesmo modelo e ao mesmo *corpus*.

6.4. Desempenho do Comitê de Classificadores

A avaliação do comitê de classificadores no contexto da **QP2** baseia-se na comparação entre o desempenho obtido por votação majoritária e o do melhor classificador individual em cada *corpus*, conforme apresentado na Tabela 5 para os *corpora* em suas versões originais.

De maneira geral, os resultados evidenciam comportamentos distintos do comitê nos diferentes *corpora* avaliados. Conforme apresentado na Tabela 5, no *corpus* ToLD-BR, observa-se um ganho de desempenho de 2% na *F-measure* para a classe tóxica, indicando que a combinação dos métodos por votação majoritária superou o melhor resultado individual. Por sua vez, o coeficiente Kappa de Cohen de 0,51 indica um nível moderado de concordância entre as predições do comitê e as anotações humanas, o que é compatível com os valores observados para os métodos individuais nesse *corpus*.

Em contrapartida, nos *corpora* HateBR e HLPHSD, o comitê não supera o método *few-shot* no cenário original, apresentando reduções marginais de *F-measure* de -2% e -1%, respectivamente (ver Tabela 5). Esse comportamento pode ser explicado pela presença de um método individual dominante, o *few-shot*, cujo desempenho superior não é compensado pela votação majoritária quando os demais métodos divergem sistematicamente de suas predições. Em cenários em que há forte disparidade de desempenho entre os métodos do comitê, a votação majoritária tende a introduzir ruído proveniente dos classificadores mais fracos, resultando em uma degradação marginal do resultado agregado (Zhou, 2012).

Em relação à **QP2**, os resultados indicam que o comitê supera o melhor método individual de forma consistente apenas no *corpus* ToLD-BR, cenário caracterizado por maior equilíbrio entre os resultados dos métodos individuais e por melhor desempenho do comitê em *F-measure* e Kappa de Cohen. Nos demais *corpora*, o comitê mantém um patamar competitivo sem sofrer degradação significativa, evidenciando robustez suficiente para operar em diferentes domínios, ainda que a estratégia de votação majoritária se mostre mais eficaz quando as métricas dos classificadores são próximas entre si. A fim de ampliar a perspectiva avaliativa, as métricas complementares de acurácia, precisão, cobertura da classe tóxica e *macro-F1*, referentes às melhores configurações dos métodos constituintes, bem como às do comitê, são apresentadas no Apêndice D (Tabela 13). Essa estratégia permite uma avaliação mais abrangente do comportamento das abordagens, sem comprometer a clareza e a objetividade da Subseção de resultados.

6.5. Impacto da Concordância Total na Anotação Automática

A avaliação do impacto da concordância total entre anotadores sobre a anotação automática, em resposta à **QP3**, baseia-se na comparação dos resultados obtidos pelos métodos nos *corpora* em suas versões originais e nos subconjuntos compostos exclusivamente por instâncias com consenso absoluto entre os avaliadores humanos. Para essa análise, adotaram-se a *F-measure* e o coeficiente Kappa de Cohen como indicadores de desempenho e de confiabilidade dos rótulos gerados.

Em razão do acentuado desbalanceamento entre as classes nos *corpora* filtrados, conforme indicado na Tabela 2, o conjunto com o maior número de instâncias foi submetido a uma subamostragem balanceada para igualar o número de exemplos tóxicos e não tóxicos. Mantêm-se, portanto, as mesmas configurações experimentais adotadas anteriormente, diferenciando-se apenas pela utilização desses subconjuntos, o que permite analisar com maior precisão o efeito das divergências entre as anotações sobre os resultados. A Tabela 6 apresenta os resultados obtidos por cada método nos *corpora* originais e filtrados, bem como as variações relativas decorrentes da filtragem. Para cada método e *corpus*, os resultados do cenário original são adotados como referência.

<i>Corpus</i>	<i>Método</i>	<i>Configuração</i>	<i>F-measure</i>	($\pm\%$)	<i>Kappa</i>	($\pm\%$)
ToLD-BR	<i>few-shot</i>	Granite3.3-8b	0.74	ref.	0.48	ref.
	Comitê		0.76	+2%	0.51	+3%
HateBR	<i>few-shot</i>	Qwen2.5-7b	0.87	ref.	0.73	ref.
	Comitê		0.85	-2%	0.72	-1%
HLPHSD	<i>few-shot</i>	Qwen2.5-7b	0.63	ref.	0.41	ref.
	Comitê		0.62	-1%	0.43	+3%

Tabela 5: Comparação entre o melhor método individual e o comitê por *corpus*, em *F-measure* e Kappa de Cohen, com variação percentual ($\pm\%$) em relação ao método individual.

<i>Corpus</i>	<i>Método</i>	<i>Original (ref.)</i>		<i>Filtrado</i>		<i>Variação</i>	
		<i>F</i>	κ	<i>F</i>	κ	$\pm\%F$	$\pm\%\kappa$
ToLD-BR	BERT	0,72	0,45	0,80	0,54	+8%	+9%
	<i>zero-shot</i>	0,73	0,47	0,80	0,55	+7%	+8%
	semi	0,72	0,46	0,79	0,56	+7%	+10%
	<i>few-shot</i>	0,74	0,48	0,81	0,56	+7%	+8%
	RAG	0,74	0,45	0,80	0,50	+6%	+5%
	Comitê	0,76	0,51	0,82	0,58	+6%	+7%
HateBR	BERT	0,83	0,66	0,86	0,74	+3%	+8%
	<i>zero-shot</i>	0,86	0,71	0,88	0,74	+2%	+3%
	semi	0,37	0,22	0,44	0,27	+7%	+5%
	<i>few-shot</i>	0,87	0,73	0,89	0,78	+2%	+5%
	RAG	0,85	0,70	0,88	0,75	+3%	+5%
	Comitê	0,85	0,72	0,89	0,79	+4%	+7%
HLPHSD	BERT	0,56	0,37	0,73	0,54	+17%	+17%
	<i>zero-shot</i>	0,62	0,38	0,82	0,61	+20%	+23%
	semi	0,41	0,06	0,38	0,23	-3%	+17%
	<i>few-shot</i>	0,63	0,41	0,83	0,64	+20%	+23%
	RAG	0,61	0,39	0,82	0,63	+21%	+24%
	Comitê	0,62	0,43	0,82	0,61	+20%	+18%

Tabela 6: *F-measure* (F) e Kappa de Cohen (κ) nos cenários original (referência) e filtrado por concordância total, com variação percentual ($\pm\%$).

De forma consistente, a filtragem por concordância total beneficia a maioria dos métodos avaliados nos três *corpora*, incluindo tanto as linhas de base *zero-shot* e BERTimbau com ajuste fino quanto os classificadores constituintes do comitê, representados pelo método semi-supervisionado baseado em grafos heterogêneos e pelas abordagens baseadas em LLMs com exemplos rotulados. Como esse padrão de melhoria é observado em métodos fundamentados em paradigmas distintos, os resultados sugerem que os ganhos obtidos estão associados à maior consistência das instâncias de referência, e não a características específicas de cada abordagem. Esse padrão também indica que parte substancial das classificações incorretas no cenário original estava relacionada a instâncias com divergência entre os anotadores, e não necessariamente a limitações intrínsecas dos métodos.

Quanto à intensidade desses ganhos, observam-se variações de modesta a expressiva conforme o *corpus* considerado. No HateBR, que já apresenta elevado nível de concordância interanotador (Kappa de Fleiss = 0,70), os métodos avaliados registram ganhos entre 2% e 7% em *F-measure*. No ToLD-BR, os ganhos mostram-se uniformes entre os métodos, situando-se entre 6% e 8%. No HLPHSD, que apresenta o menor índice de concordância interanotador (Kappa de Fleiss = 0,17 na tarefa binária), os métodos alcançam ganhos de até 21% em *F-measure* e de 24% em Kappa de Cohen.

No que se refere ao método semi-supervisionado, observam-se ganhos de desempenho nos *corpora* ToLD-BR e HateBR após a filtragem. No ToLD-BR, esses ganhos são mais evidentes, enquanto, no HateBR, o método permanece aquém das abordagens baseadas em LLMs. No *corpus* HLPHSD, por sua vez, o método registra uma leve queda na

F-measure (-3%), apesar de um discreto ganho no coeficiente Kappa de Cohen, que mantém um nível razoável de concordância com as anotações humanas. Este desempenho do método semi-supervisionado nos três *corpora* reflete um cenário que combina mudança de domínio e redução do volume de dados disponíveis após a filtragem e a subamostragem balanceada. Essas condições relacionam-se à sensibilidade do método à estrutura do grafo heterogêneo, que depende da coocorrência de termos entre instâncias rotuladas e não rotuladas. Nesse contexto, a redução do conjunto de dados e as diferenças de domínio tendem a alterar esses padrões de coocorrência, o que se reflete em variações na estrutura do grafo em cenários de menor desempenho do método.

De maneira geral, esses resultados indicam que a filtragem por concordância total eleva o desempenho dos métodos individuais a um patamar já bastante consolidado, no qual os ganhos adicionais em *F-measure* obtidos pela agregação por votação majoritária tornam-se naturalmente mais discretos. Cumpre destacar, contudo, que o comitê apresenta o maior Kappa de Cohen no cenário filtrado nos *corpora* ToLD-BR e HateBR, mantendo-se em patamar competitivo no HLPDSD, evidenciando que a agregação contribui para um alinhamento consistente com as anotações humanas, mesmo em condições em que os métodos individuais já apresentam desempenho elevado.

No que se refere à **QP3**, os resultados indicam que todos os métodos avaliados, incluindo o comitê, apresentam desempenho consistentemente superior quando aplicados a instâncias com concordância total entre os anotadores, mantendo as mesmas configurações experimentais adotadas no cenário original. Ainda que as diferenças entre os *corpora* possam decorrer de múltiplos fatores, como protocolos de anotação, diretrizes adotadas e perfil dos anotadores, os ganhos observados após a filtragem evidenciam que a consistência das anotações de referência desempenha papel central na avaliação de métodos automáticos de identificação de linguagem tóxica, tarefa marcada por elevado grau de subjetividade. Sob as mesmas condições experimentais, subconjuntos compostos por instâncias com maior consenso entre anotadores associam-se a melhores níveis de desempenho em todos os paradigmas avaliados.

6.6. Análise Qualitativa dos Erros e Avaliação de Equidade

Esta Subseção apresenta uma análise qualitativa dos resultados obtidos nos dois cenários experimentais, com ênfase no desempenho do comitê. Adicionalmente, avalia-se o grau de viés associado ao conteúdo tóxico por meio de uma lista pré-definida de termos de identidade, adaptada de (Dixon et al., 2018; Kennedy et al., 2020), que foi traduzida e ampliada manualmente para contemplar as especificidades do contexto socio-cultural brasileiro (ver Apêndice C). Para tanto, a Tabela 7 apresenta as matrizes de confusão do comitê para os *corpora* avaliados, permitindo identificar os padrões de erro em cada cenário. Nesse contexto, a distinção entre falsos positivos e falsos negativos assume um papel central. Falsos positivos correspondem à atribuição indevida do rótulo de toxicidade a textos não tóxicos, ou seja, à interpretação equivocada de ofensa em conteúdos legítimos; por sua vez, falsos negativos referem-se à não identificação de instâncias efetivamente tóxicas, permitindo que conteúdos potencialmente prejudiciais permaneçam sem rotulação adequada. Em aplicações de detecção automática de toxicidade, esse tipo de erro tende a ser particularmente crítico, pois permite que conteúdos ofensivos ou prejudiciais continuem circulando nas plataformas sem qualquer sinalização ou intervenção. A fim de analisar os padrões de erro, as matrizes de confusão apresentam, além das contagens absolutas, percentuais normalizados por linha (classe verdadeira), permitindo observar a proporção relativa de acertos e erros em cada classe.

De forma complementar, a Tabela 8 apresenta exemplos textuais extraídos dos *corpora* originais, acompanhados dos rótulos atribuídos pelos anotadores humanos, das predições individuais dos métodos que compõem o comitê e da decisão final agregada. Nessa tabela, o valor **1** corresponde à classe **tóxico**, enquanto o valor **0** corresponde à classe **não tóxico**. A análise desses exemplos contribui para uma compreensão mais aprofundada das fragilidades ainda presentes no comitê, particularmente em casos que envolvem toxicidade implícita, ambiguidade contextual ou uso irônico da linguagem.

No *corpus* ToLD-BR, observa-se que o comitê apresenta, no cenário original, um número expressivo de erros de classificação, com predominância de falsos positivos, que correspondem a 36,4% das instâncias não tóxicas classificadas incorretamente como tóxicas, enquanto os falsos negativos representam 11,3% das instâncias tóxicas não identificadas pelo modelo, con-

Classe Verdadeira	Classe Predita	
	Tóxico	Não tóxico
Tóxico	8213 (88.7%)	1042 (11.3%)
Não tóxico	4273 (36.4%)	7472 (63.6%)

(a) comitê – ToLD-BR (original)

Classe Verdadeira	Classe Predita	
	Tóxico	Não tóxico
Tóxico	2863 (81.8%)	637 (18.2%)
Não tóxico	358 (11.1%)	2863 (88.9%)

(c) comitê – HateBR (original)

Classe Verdadeira	Classe Predita	
	Tóxico	Não tóxico
Tóxico	1220 (68.2%)	568 (31.8%)
Não tóxico	903 (23.3%)	2979 (76.7%)

(e) comitê – HLPDSD (original)

Classe Verdadeira	Classe Predita	
	Tóxico	Não tóxico
Tóxico	1380 (95.0%)	73 (5.0%)
Não tóxico	554 (38.1%)	899 (61.9%)

(b) comitê – ToLD-BR (filtrado)

Classe Verdadeira	Classe Predita	
	Tóxico	Não tóxico
Tóxico	2092 (87.8%)	290 (12.2%)
Não tóxico	207 (8.7%)	2175 (91.3%)

(d) comitê – HateBR (filtrado)

Classe Verdadeira	Classe Predita	
	Tóxico	Não tóxico
Tóxico	421 (83.0%)	86 (17.0%)
Não tóxico	87 (17.2%)	420 (82.8%)

(f) comitê – HLPDSD (filtrado)

Tabela 7: Matrizes de confusão do comitê de classificadores para os *corpora* ToLD-BR, HateBR e HLPDSD, considerando os conjuntos em sua versão original e após a filtragem por concordância total entre anotadores. Os valores indicam contagens absolutas, com percentuais normalizados por linha (classe verdadeira) apresentados entre parênteses.

forme evidenciado na matriz de confusão da Tabela 7(a). Esse padrão indica maior sensibilidade dos classificadores do comitê à classe tóxica, em detrimento da especificidade na identificação de conteúdos não tóxicos. No cenário filtrado por concordância total (ver Tabela 7(b)), verifica-se uma redução substancial dos erros, com apenas 554 falsos positivos (38,1%) e 73 falsos negativos (5,0%), evidenciando o impacto positivo da remoção de instâncias com divergência anotativa. Ainda assim, a persistência de um número elevado de falsos positivos sugere que a dificuldade não decorre exclusivamente das divergências interpretativas na anotação, mas também da identificação inadequada de termos ofensivos empregados em contextos não agressivos, como em situações de autodepreciação ou humor. Nesses casos, os classificadores tendem a priorizar o sinal lexical da palavra ofensiva, em detrimento da interpretação pragmática do texto. Tal situação pode ser observada no exemplo iv) da Tabela 8, “nooo mano kkk eu sou muito burro to bolado”, em que o termo “burro” conduz todos os métodos que compõem o comitê a uma interpretação literal de toxicidade. Por outro lado, os falsos negativos tendem a ocorrer em sentenças nas quais a ofensa é veiculada de forma implícita ou metafórica, exigindo inferência pragmática para a correta identificação da toxicidade, como no exemplo ii) “eu pensava que era uma vaca a burra meu pai”, em que o insulto é direcionado a um

terceiro sem o uso de marcadores convencionais de discurso ofensivo.

Em contraste com o padrão observado no *corpus* ToLD-BR, em que os falsos positivos predominam, o *corpus* HateBR apresenta uma inversão na distribuição dos erros. Conforme a matriz de confusão do comitê (Tabela 7(c)), observa-se maior incidência de falsos negativos (637), o que corresponde a 18,2% das instâncias tóxicas não identificadas pelo comitê. Ainda assim, a proporção geral de erros é inferior à observada no ToLD-BR, o que indica que o comitê identifica com maior clareza a manifestação da toxicidade no HateBR. Quando consideradas apenas as instâncias com concordância total entre os anotadores (Tabela 7(d)), observa-se uma redução consistente em ambos os tipos de erro, com 12,2% de falsos negativos e 8,7% de falsos positivos, reforçando o efeito positivo da remoção de instâncias com divergência nos rótulos. No plano qualitativo, os exemplos (vi) e (vii) da Tabela 8 ilustram erros associados ao emprego de ironia e sarcasmo, em que textos rotulados como não tóxicos são interpretados pelo comitê como ataques ou desqualificações, o que contribui para a ocorrência de falsos positivos. Já os exemplos viii) e x) evidenciam a situação inversa, em que a toxicidade é expressa de forma indireta e depende de conhecimento contextual, o que favorece falsos negativos quando o teor ofensivo não se manifesta de forma lexicalmente explícita.

Corpus	Exemplo textual	Rótulo original	semi	few	RAG	Comitê
ToLD-BR	i) sou a favor de invadir o stf e não chegarmos lá de mãos vazias....	1	0	1	1	1
	ii) eu pensava que era uma vaca a burra meu pai	1	1	0	0	0
	iii) gente eu não sou feia parem de me tratar assim	0	0	0	1	0
	iv) nooo mano kkk eu sou muito burro to bolado	0	1	1	1	1
	v) eu também tiro foto com aviados e simpatizo mas isso não quer dizer nada se eu pudesse ainda metia uma surra	1	0	1	1	1
HateBR	vi) jesus é cada uma desse cidadão é de assustar	0	0	1	1	1
	vii) é sério que estão comemorando a invasão dos nazistas caralho decepção	0	1	1	1	1
	viii) mais verdadeira que nota de 3 reais	1	0	1	0	0
	ix) hipócrita	1	0	1	1	1
	x) faz um programa com o lula. sugestão para o nome: sol quadrado	1	0	0	0	0
HLPHSD	xi) só matando amiga	1	0	1	1	1
	xii) falou populista já sei que é da turma fake news e retardado	0	1	1	1	1
	xiii) tem mão de homem ou seja com calos	1	1	0	0	0
	xiv) os “refugiados” e a classe política traidora	0	1	1	1	1
	xv) abri a janela tava uma gorda no apartamento da frente olhando pra baixo e chorando acho que caiu a empadinha dela	1	0	0	0	0

Tabela 8: Exemplos textuais com rótulos humanos, predições dos métodos individuais (semi, few, RAG) e decisão do comitê. Os valores 1 e 0 indicam classe tóxico e não tóxico, respectivamente.

No *corpus* HLPHSD, o padrão de erro no cenário original revela predominância de falsos positivos (903) em relação aos falsos negativos (568), conforme indicado na Tabela 7(e). Ao se considerar o subconjunto filtrado por concordância total entre anotadores (Tabela 7(f)), observa-se uma redução expressiva e praticamente equilibrada dos erros, com 86 falsos negativos e 87 falsos positivos, sugerindo que parcela significativa das classificações incorretas no cenário original pode estar relacionada a instâncias com divergência interanotador. Em termos qualitativos, o exemplo xiii) “*tem mão de homem ou seja com calos*” ilustra um caso de toxicidade sutil, na medida em que pressupõe que a masculinidade estaria condicionada à presença de determinados traços físicos, como calos nas mãos, implicando que a ausência desses traços descaracterizaria o sujeito como homem. Por outro lado, os exemplos xii) e xiv) indicam possíveis inconsistências na rotulagem manual. No primeiro caso, verifica-se uma ofensa direta, na medida em que o termo pejorativo retardado é empregado para qualificar o interlocutor no contexto do debate em que o texto foi coletado, o que justificaria sua classificação como tóxico. Já no segundo, a associação do substantivo entre aspas “refugiados” à figura de traidor pode igualmente configurar um ataque ou uma

desqualificação implícita. Em ambos os casos, o comitê classificou os textos como tóxicos, sugerindo que parte dos erros atribuídos ao sistema reflete divergências na rotulagem realizada pelos anotadores.

Para além da análise binária dos textos tóxicos, é importante observar que determinados termos ou expressões podem representar formas específicas de toxicidade direcionadas a indivíduos ou grupos com base em atributos como aparência física, orientação sexual ou raça, o que pode introduzir desafios adicionais para os classificadores na identificação desse fenômeno. No exemplo xv), “*abri a janela tava uma gorda no apartamento da frente olhando pra baixo e chorando acho que caiu a empadinha dela*”, a ofensa dirige-se à aparência física, configurando um caso de gordofobia. Mesmo quando a intenção depreciativa pode ser inferida do contexto, os classificadores não conseguem associar o termo “gorda” a essa manifestação específica de toxicidade, classificando o texto como não tóxico. De maneira semelhante, o exemplo v) “*eu também tiro foto com aviados e simpatizo mas isso não quer dizer nada se eu pudesse ainda metia uma surra*” apresenta uma variação lexical (“aviados”) potencialmente empregada para camuflar a ofensa. Nesse caso, os classificadores baseados em LLMs conseguem interpretar o contexto discursivo e inferir

que a toxicidade está associada a um ataque de natureza homofóbica, enquanto o método semi-supervisionado falha nessa detecção. Termos com baixa conectividade no grafo heterogêneo, como a variação lexical “aviados”, dispõem de menos evidências para a propagação de rótulos, conforme já discutido na Subseção 6.5.

Para complementar a análise qualitativa com uma avaliação quantitativa de disparidades, empregaram-se métricas de avaliação de equidade (Borkan et al., 2019), combinando três métricas baseadas na área sob a curva característica de operação, conhecida como AUC (do inglês *Area Under the Curve*): (i) *Subgroup AUC*, que mensura a capacidade do modelo de discriminar entre tóxico e não tóxico dentro do conjunto de instâncias identitárias; (ii) *Background Positive, Subgroup Negative* (BPSN), que avalia a sensibilidade a falsos positivos no grupo, isto é, textos não tóxicos do grupo classificados como tóxicos por associação com marcadores identitários; e (iii) *Background Negative, Subgroup Positive* (BNSP), que avalia a sensibilidade a falsos negativos, ou seja, textos tóxicos do grupo não identificados pelo modelo. Em todas essas métricas, valores mais próximos de 1 indicam menor disparidade. A análise foi conduzida de forma agregada, considerando todas as instâncias que contêm ao menos um termo da lista identitária como um único grupo (ver Apêndice C), em coerência com o tratamento binário de toxicidade adotado em todo o trabalho.

Corpus	Método	n	Sub. AUC	BPSN	BNSP
ToLD-BR	semi	1560	0,67	0,66	0,76
	<i>few-shot</i>	1560	0,68	0,69	0,75
	RAG	1560	0,68	0,68	0,73
HateBR	semi	570	0,62	0,61	0,64
	<i>few-shot</i>	570	0,79	0,77	0,88
	RAG	570	0,78	0,77	0,87
HLPHSD	semi	1251	0,63	0,51	0,68
	<i>few-shot</i>	1251	0,70	0,68	0,75
	RAG	1251	0,68	0,67	0,73

Tabela 9: Resultados das métricas de equidade (*Subgroup AUC*, BPSN e BNSP) para as instâncias contendo termos identitários nos *corpora* avaliados. Valores próximos de 1 indicam maior equidade e ausência de viés entre instâncias identitárias e não identitárias.

A análise das métricas apresentadas na Tabela 9 revela comportamentos distintos entre os classificadores do comitê e os *corpora* avaliados. No ToLD-BR, os três classificadores apresentam métricas uniformes, com *Subgroup AUC* em torno de 0,68, BPSN entre 0,66 e 0,69 e BNSP entre 0,73 e 0,76. Os valores de BNSP, ligeiramente

superiores aos de BPSN, indicam maior estabilidade dos métodos na identificação de toxicidade efetiva em instâncias com termos identitários do que na distinção entre textos não tóxicos do grupo e textos tóxicos do restante do *corpus*. Nos *corpora* HateBR e HLPHSD, esse padrão uniforme dá lugar a uma discrepância acentuada entre os classificadores. Os métodos baseados em LLMs se mantêm em patamares elevados, com *Subgroup AUC* entre 0,68 e 0,79, BPSN entre 0,67 e 0,77 e BNSP entre 0,73 e 0,88. O valor elevado de BNSP sugere baixa propensão a falsos negativos em instâncias identitárias. O método semi-supervisionado, em contrapartida, registra valores substancialmente inferiores, com destaque para o BPSN de 0,51 no HLPHSD, próximo ao desempenho aleatório. Esse desequilíbrio entre as métricas BPSN e BNSP no método semi-supervisionado revela que o classificador aprende a associação entre os marcadores de identidade e os rótulos tóxicos no *corpus* de treinamento, configurando um caso típico de correlação espúria, em vez de aprender propriedades semânticas da toxicidade em si.

De modo geral, as análises realizadas nesta Subseção indicam que o comitê de classificadores enfrenta dificuldades na identificação de toxicidade implícita, como ironia, sarcasmo e ofensas dependentes de contexto, que frequentemente demandam conhecimento adicional para serem corretamente interpretadas. A presença de termos comumente associados a contextos ofensivos (e.g., “burro”, “nazista”) contribui para o aumento da taxa de falsos positivos, na medida em que os métodos tendem a supervalorizar termos isolados em detrimento do significado contextual. Ademais, as métricas de equidade mostram que vieses presentes nos dados de treinamento podem se propagar aos rótulos gerados automaticamente. Esse achado é particularmente relevante no contexto deste trabalho, cuja contribuição central reside em uma metodologia para anotação automática de *corpora* de linguagem tóxica, uma vez que potenciais vieses identificados na fase de anotação podem se propagar aos recursos linguísticos produzidos por meio dessa estratégia e às aplicações que venham a utilizá-los. A análise de equidade conduzida nesta Subseção contribui justamente para mapear esse problema, fornecendo subsídios para o uso responsável da abordagem proposta na construção de *corpora* anotados. A mitigação dessas disparidades por meio de estratégias de *debiasing*, balanceamento de instâncias identitárias no *corpus* de treinamento e análise de sensibilidade a *prompts* constitui uma direção prioritária para trabalhos futuros (Navigli et al., 2023; Dixon et al., 2018).

6.7. Concordância entre os Métodos do Comitê

Com o objetivo de aprofundar a avaliação do comitê de classificadores, esta seção analisa o grau de concordância entre os métodos constituintes com base em suas classificações nos *corpora* avaliados. A Tabela 10 apresenta os valores do coeficiente Kappa de Cohen calculados para cada par de métodos, considerando os cenários original e filtrado por concordância total. Essa métrica permite quantificar o grau de alinhamento entre as decisões dos classificadores e fornece uma estimativa mais robusta da similaridade entre suas predições, complementando a avaliação de desempenho apresentada nas Subseções anteriores.

Cenário	Corpus	semi vs <i>few</i>	semi vs RAG	<i>few</i> vs RAG
Original	ToLD-BR	0.66	0.63	0.72
	HateBR	0.24	0.24	0.84
	HLPHSD	0.20	0.20	0.76
Filtrado	ToLD-BR	0.69	0.61	0.82
	HateBR	0.24	0.23	0.77
	HLPHSD	0.24	0.21	0.82

Tabela 10: Concordância entre os métodos constituintes do comitê, mensurada pelo coeficiente Kappa de Cohen para cada par de abordagens, nos cenários original e filtrado por concordância total.

Considerando os *corpora* em sua versão original, no ToLD-BR observa-se concordância substancial entre todos os pares de métodos (Kappa entre 0,63 e 0,72), padrão consistente com o desempenho equilibrado relatado na Subseção 6.2, o que ajuda a explicar os ganhos do comitê nesse corpus por meio de votação majoritária.

Nos *corpora* HateBR e HLPHSD, as diferenças de desempenho entre os métodos tornam-se mais acentuadas, acompanhadas de uma mudança nos valores de concordância entre os pares de classificadores. A concordância entre os métodos baseados em LLMs permanece elevada, com Kappa de 0,84 no HateBR e de 0,76 no HLPHSD, indicando forte alinhamento entre os rótulos gerados por essas abordagens. Por sua vez, a concordância no método semi-supervisionado é consideravelmente inferior, situando-se entre 0,20 e 0,24 em ambos os *corpora*. Esse contraste reforça que a métrica de concordância entre os métodos acompanha o comportamento de desempenho observado nessas *corpora*. As abordagens baseadas em LLMs apresentam maior alinhamento entre os rótulos gerados e os melhores resultados em *F-measure*,

enquanto o método semi-supervisionado, cujos rótulos divergem dos demais, apresenta também o desempenho mais baixo. Essa observação indica que a votação majoritária tende a refletir o alinhamento do subgrupo mais coeso, que, nesses *corpora*, corresponde às abordagens baseadas em LLMs, com os ganhos do comitê limitados pela divergência nos rótulos do método semi-supervisionado.

No cenário filtrado por concordância total, as observações permanecem, em linhas gerais, consistentes com as do cenário original. A concordância entre o método semi-supervisionado e as abordagens baseadas em LLMs mantém-se substancial apenas no ToLD-BR, permanecendo fraca nos demais *corpora*. Um aspecto relevante, contudo, é o aumento no alinhamento entre os métodos baseados em LLMs após a filtragem. No ToLD-BR, o coeficiente Kappa entre *few-shot* e RAG evoluiu de 0,72 para 0,82, indicando concordância quase perfeita. Uma tendência semelhante é observada no HLPHSD, em que o valor passa de 0,76 para 0,82.

Quanto à **QP4**, os resultados indicam que o grau de concordância entre os classificadores do comitê está positivamente associado aos ganhos obtidos pela estratégia de votação majoritária. No ToLD-BR, em que se observa concordância substancial entre todos os pares de métodos, o comitê obteve ganhos consistentes em relação ao melhor método individual. Nos *corpora* HateBR e HLPHSD, a forte divergência entre o método semi-supervisionado e as abordagens baseadas em LLMs limitou os benefícios da agregação, resultando em desempenho equiparável ou marginalmente inferior ao do melhor classificador isolado. Em conjunto, esses achados sugerem que a eficácia da votação majoritária depende não apenas da qualidade individual dos métodos combinados, mas também do grau de alinhamento entre os rótulos por eles gerados. Nesse sentido, a investigação de estratégias de agregação adaptativas, como votação ponderada ou *stacking*, que incorporem o grau de concordância entre os classificadores constituintes como critério de ponderação, representa um desdobramento natural desta investigação.

7. Conclusão

Este trabalho investigou o uso de um comitê de classificadores para a anotação automática de linguagem tóxica em português, considerando cenários de escassez de dados rotulados e de mudança de domínio. A proposta integrada três abordagens complementares funda-

mentadas em paradigmas distintos: um método semi-supervisionado baseado em grafos heterogêneos, uma estratégia de inferência *few-shot* com grandes modelos de linguagem e um método baseado em *Retrieval-Augmented Generation*. Essas abordagens foram combinadas por meio de votação majoritária e avaliadas em múltiplos *corpora*, abrangendo diferentes domínios e plataformas de coleta.

De forma consolidada, os resultados indicam que o comitê proposto constitui uma estratégia metodologicamente consistente e adequada para a anotação automática de linguagem tóxica. O comitê supera os classificadores de linha de base, *zero-shot* e BERTimbau com ajuste fino, na maioria dos cenários avaliados, e mantém desempenho competitivo mesmo quando, isoladamente, não supera o melhor método individual. Em cenários caracterizados por maior equilíbrio de desempenho e alinhamento entre os métodos constituintes, como observado no ToLD-BR, o comitê obtém ganhos consistentes, com melhoria de até 2% na *F-measure* e aumento dos coeficientes Kappa de Cohen. Em *corpora* caracterizados por maior divergência entre os métodos constituintes do comitê, como observado no HateBR e no HLPHSD, o desempenho do comitê mantém-se próximo ao do melhor método individual.

A análise comparativa entre os cenários experimentais, considerando os *corpora* em sua versão original e após a filtragem por concordância total, evidencia que o comitê apresenta o maior coeficiente Kappa de Cohen no cenário filtrado nos *corpora* ToLD-BR e HateBR, mantendo-se em patamar competitivo no HLPHSD, o que sugere que a votação majoritária se beneficia de decisões individuais mais alinhadas. Cumpre destacar que tais resultados foram obtidos a partir de um único *corpus* curado, composto por 1.400 instâncias, utilizado como base para treinamento e inferência, enquanto a avaliação do comitê foi conduzida em conjuntos substancialmente maiores, provenientes de diferentes domínios e plataformas. Esse delineamento experimental indica o potencial da abordagem para ampliar, de forma escalável, os recursos anotados, mesmo em contextos com limitada disponibilidade de dados rotulados.

Complementarmente, a análise qualitativa dos erros e a avaliação de equidade por grupo identitário evidenciaram que o comitê ainda enfrenta dificuldades na identificação de toxicidade implícita, como ironia, sarcasmo e ofensas dependentes de contexto, e que vieses presentes nos dados de treinamento podem se propagar aos

rótulos gerados automaticamente. Esse achado é particularmente relevante no contexto deste trabalho, cuja contribuição central reside em uma metodologia para construção ou ampliação de *corpora* anotados, uma vez que dados com vieses comprometem não apenas a qualidade dos recursos produzidos, mas também a das aplicações que venham a utilizá-los.

Entre as limitações deste estudo, destacam-se a utilização de um único *corpus* curado como base de treinamento, o emprego de modelos multilíngues de pequeno porte, sem ajuste fino específico ao domínio, e a ausência de mecanismos de interpretabilidade que permitam explicar, de forma sistemática, as decisões individuais dos classificadores que compõem o comitê. Como trabalhos futuros desta investigação, pretende-se ampliar o treinamento do comitê para múltiplos *corpora*, diversificando os domínios textuais utilizados como base e explorando o uso de modelos de linguagem de maior porte, potencialmente mais aptos a apreender nuances semânticas e subjetivas. Adicionalmente, a incorporação de outras estratégias de agregação, como votação ponderada e *stacking*, bem como o desenvolvimento de mecanismos sensíveis ao contexto discursivo e de adaptação de domínio (Jiang & Zubiaga, 2023; Safikhani & Broneske, 2025), configuram uma direção promissora para o aprimoramento da robustez e da capacidade de generalização da metodologia proposta. No que concerne à avaliação de equidade, a investigação de estratégias de mitigação de vieses, o balanceamento de instâncias identitárias no *corpus* de treinamento e a análise de sensibilidade a *prompts* constituem direções prioritárias (Navigli et al., 2023; Dixon et al., 2018). Por fim, a incorporação de técnicas de interpretabilidade *post-hoc* (Ribeiro et al., 2016), com avaliação da fidelidade das explicações (Jacovi & Goldberg, 2020; Salles et al., 2025), permitirá explicitar os critérios subjacentes às decisões dos classificadores e contribuirá para a construção de *corpora* anotados com maior transparência metodológica.

Referências

- Alsafari, Safa & Samira Sadaoui. 2021. Semi-supervised self-training of hate and offensive speech from social media. *Applied Artificial Intelligence* 35(15). 1621–1645. doi 10.1080/08839514.2021.1988443
- Aroyo, Lora, Lucas Dixon, Nithum Thain, Olivia Redfield & Rachel Rosen. 2019. Crowdsourcing subjective tasks: The case study of understanding toxicity in online discus-

- sions. Em *Companion Proceedings of The 2019 World Wide Web Conference*, 1100–1105. doi 10.1145/3308560.3317083
- Assis, Gabriel, Annie Amorim, Jonnathan Carvalho, Mariza Ferro, Daniel de Oliveira, Daniela Vianna & Aline Paes. 2024. Explorando técnicas de aprendizado em modelos de linguagem para classificação de discurso de Ódio e ofensivo em Português. *Linguamática* 16(2). 91–113. doi 10.21814/lm.16.2.446
- Badjatiya, Pinkesh, Shashank Gupta, Manish Gupta & Vasudeva Varma. 2017. Deep learning for hate speech detection in Tweets. Em *26th International Conference on World Wide Web Companion*, 759–760. doi 10.1145/3041021.3054223
- Borkan, Daniel, Lucas Dixon, Jeffrey Sorensen, Nithum Thain & Lucy Vasserman. 2019. Nuanced metrics for measuring unintended bias with real data for text classification. Em *World Wide Web Conference (WWW)*, 491–500. doi 10.1145/3308560.3317593
- Botella-Gil, Beatriz, Robiert Sepúlveda-Torres, Alba Bonet-Jover, Patricio Martínez-Barco & Estela Saquete. 2024. Semi-automatic dataset annotation applied to automatic violent message detection. *IEEE Access* 12. 19651–19664. doi 10.1109/ACCESS.2024.3361404
- Brasil, Presidência da República. 1989. Lei nº 7.716, de 5 de janeiro de 1989. Acesso em: jul. 2025. [↗](#)
- Brown, Tom, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D 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 Ziegler, Jeffrey Wu, Clemens Winter, Chris Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever & Dario Amodei. 2020. Language models are few-shot learners. Em *Advances in Neural Information Processing Systems (NeurIPS)*, 1877–1901. [↗](#)
- Butler, Judith. 2021. *Discurso de ódio: Uma política do performativo*. São Paulo: Editora Unesp
- Cjadam, Jeffrey Sorensen, Julia Elliott, Lucas Dixon, Mark McDonald, Nithum & Will Cukierski. 2017. Toxic comment classification challenge. Kaggle. [↗](#)
- Costa Bertaglia, Thales Felipe & Maria das Graças Volpe Nunes. 2016. Exploring word embeddings for unsupervised textual user-generated content normalization. Em *2nd Workshop on Noisy User-generated Text (WNUT)*, 112–120. [↗](#)
- Devlin, Jacob, Ming-Wei Chang, Kenton Lee & Kristina Toutanova. 2019. BERT: Pre-training of deep bidirectional transformers for language understanding. Em *Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL)*, 4171–4186. doi 10.18653/v1/N19-1423
- Dirting, Bakwa Dunka, Gloria A. Chukwudebe, Euphemia Chioma Nwokorie & Ikechukwu Ignatius Ayogu. 2022. Multi-label classification of hate speech severity on social media using BERT model. Em *4th International Conference on Disruptive Technologies for Sustainable Development (NIGERCON)*, 1–5. doi 10.1109/NIGERCON54645.2022.9803164
- Dixon, Lucas, John Li, Jeffrey Sorensen, Nithum Thain & Lucy Vasserman. 2018. Measuring and mitigating unintended bias in text classification. Em *AAAI/ACM Conference on AI, Ethics, and Society*, 67–73. doi 10.1145/3278721.3278729
- Fortuna, Paula & Sérgio Nunes. 2018. A survey on automatic detection of hate speech in text. *ACM Computing Survey* 51(4). 1–30. doi 10.1145/3232676
- Fortuna, Paula, João Rocha da Silva, Juan Soler-Company, Leo Wanner & Sérgio Nunes. 2019. A hierarchically-labeled Portuguese hate speech dataset. Em *3rd Workshop on Abusive Language Online*, 94–104. doi 10.18653/v1/W19-3510
- Founta, Antogni, Constantinos Djouvas, Despoina Chatzakou, Ilias Leontiadis, Jeremy Blackburn, Gianluca Stringhini, Athena Vakali, Michael Sirivianos & Nicolas Kourtellis. 2018. Large scale crowdsourcing and characterization of Twitter abusive behavior. Em *International AAAI Conference on Web and Social Media (ICWSM)*, doi 10.1609/icwsml.v12i1.14991
- Frediani, João Otávio Rodrigues Ferreira, Gabriel Lino Garcia, Pedro Henrique Paiola, Leandro Aparecido Passos, João Paulo Papa & Aparecido Nilceu Marana. 2025. Hate speech detection in Portuguese using BERTimbau. Em *Progress in Pattern Recognition*,

- Image Analysis, Computer Vision, and Applications*, 244–255
- Gilardi, Fabrizio, Meysam Alizadeh & Maël Kubli. 2023. ChatGPT outperforms crowd workers for text-annotation tasks. *Proceedings of the National Academy of Sciences* 120(30). doi 10.1073/pnas.2305016120
- Hettiachchi, Danula, Indigo Holcombe-James, Stephanie Livingstone, Anjalee de Silva, Matthew Lease, Flora Salim & Mark Sander-son. 2023. How crowd worker factors influence subjective annotations: A study of tagging misogynistic hate speech in tweets. Em *Conference on Human Computation and Crowdsourcing*, 38–50. doi 10.1609/hcomp.v11i1.27546
- IBM Research. 2025. IBM Granite-3.3 language models GitHub repository. Apache-2.0 licensed repository for Granite 3.3 language models, including 8B and 2B variants. Accessed: 2026-02-18. [↗](#)
- Jacovi, Alon & Yoav Goldberg. 2020. Towards faithfully interpretable NLP systems: How should we define and evaluate faithfulness? Em *58th Annual Meeting of the Association for Computational Linguistics (ACL)*, 4198–4205. doi 10.18653/v1/2020.acl-main.386
- Jahan, Md Saroar & Mourad Oussalah. 2023. A systematic review of hate speech automatic detection using natural language processing. *Neurocomputing* 546. 126232. doi 10.1016/j.neucom.2023.126232
- Jiang, Aiqi & Arkaitz Zubiaga. 2023. SexWEs: Domain-aware word embeddings via cross-lingual semantic specialisation for chinese sexism detection in social media. Em *Proceedings of the International AAAI Conference on Web and Social Media (ICWSM)*, 447–458. doi 10.1609/icwsml.v17i1.22159
- Kennedy, Brendan, Xisen Jin, Aida Mostafazadeh Davani, Morteza Dehghani & Xiang Ren. 2020. Contextualizing hate speech classifiers with post-hoc explanation. Em *58th Annual Meeting of the Association for Computational Linguistics (ACL)*, 5435–5442. doi 10.18653/v1/2020.acl-main.483
- King, Ben, Rahul Jha & Dragomir R. Radev. 2014. Heterogeneous networks and their applications: Scientometrics, name disambiguation, and topic modeling. *Transactions of the Association for Computational Linguistics (TACL)* 2. 1–14. doi 10.1162/tac1_a_00161
- Landis, J. Richard & Gary G. Koch. 1977. The measurement of observer agreement for categorical data. *Biometrics* 33(1). 159–174. doi 10.2307/2529310
- Lees, Alyssa, Vinh Q. Tran, Yi Tay, Jeffrey Sorensen, Jai Gupta, Donald Metzler & Lucy Vasserman. 2022. A new generation of perspective api: Efficient multilingual character-level transformers. Em *28th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, 3197–3207. doi 10.1145/3534678.3539147
- Leite, João Augusto, Diego Silva, Kalina Bontcheva & Carolina Scarton. 2020. Toxic language detection in social media for Brazilian Portuguese: New dataset and multilingual analysis. Em *1st Conference of the Asia-Pacific Chapter of the Association for Computational Linguistics (AACL) and the 10th International Joint Conference on Natural Language Processing*, 914–924. doi 10.18653/v1/2020.aacl-main.91
- Lewis, Patrick, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wentaou Yih, Tim Rocktäschel, Sebastian Riedel & Douwe Kiela. 2020. Retrieval-augmented generation for knowledge-intensive NLP tasks. Em *34th International Conference on Neural Information Processing Systems (NIPS)*, [↗](#)
- Mladenović, Miljana, Vera Ošmjanski & Staša Vujičić Stanković. 2021. Cyber-aggression, cyberbullying, and cyber-grooming: A survey and research challenges. *ACM Computing Surveys* 54(1). 1–42. doi 10.1145/3424246
- Nascimento, Gabriel, Flavio Carvalho, Alexandre Martins da Cunha, Carlos Roberto Viana & Gustavo Paiva Guedes. 2019. Hate speech detection using brazilian imageboards. Em *25th Brazilian Symposium on Multimedia and the Web*, 325–328. doi 10.1145/3323503.3360619
- Navigli, Roberto, Simone Conia & Björn Ross. 2023. Biases in large language models: Origins, inventory, and discussion. *Journal of Data and Information Quality* 15(2). 1–21. doi 10.1145/3597307
- Neto, Francisco R., Rafael Anchiêta, Raimundo Moura & André Santana. 2024. Abordagem semi-supervisionada para anotação de linguagem tóxica. Em *Anais do XIII Brazilian Workshop on Social Network Analysis and Mining (BRASNAM)*, 116–129. doi 10.5753/brasnam.2024.2965
- Oliveira, Amanda, Thiago Cecote, Pedro Silva, Jadson Gertrudes, Vander Freitas & Eduardo

- Luz. 2023. How good is ChatGPT for detecting hate speech in Portuguese? Em *Anais do XIV Simpósio Brasileiro de Tecnologia da Informação e da Linguagem Humana (STIL)*, 103–112. [↗](#)
- Oliveira, Amanda, Pedro H. Silva, Valéria Santos, Gladston Moreira, Vander L. Freitas & Eduardo J. Luz. 2024. Toxic text classification in Portuguese: Is LLaMA 3.1 8B all you need? Em *Anais do XV Simpósio Brasileiro de Tecnologia da Informação e da Linguagem Humana (STIL)*, 57–66. [↗](#)
- Pedregosa, F., G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, J. Vanderplas, A. Passos, D. Cournapeau, M. Brucher, M. Perrot & E. Duchesnay. 2011. Scikit-learn: Machine learning in Python. *Journal of Machine Learning Research* 12. 2825–2830. [↗](#)
- Pelle, Rogers, Cleber Alcântara & Viviane P. Moreira. 2018. A classifier ensemble for offensive text detection. Em *24th Brazilian Symposium on Multimedia and the Web*, 237–243. [doi](#) 10.1145/3243082.3243111
- de Pelle, Rogers Prates & Viviane P. Moreira. 2017. Offensive comments in the Brazilian Web: a dataset and baseline results. Em *Anais do VI Brazilian Workshop on Social Network Analysis and Mining (BRASNAM)*, 510–519. [doi](#) 10.5753/brasnam.2017.3260
- Pelosi, Serena, Alessandro Maisto, Pierluigi Vitale & Simonetta Vietri. 2017. Mining offensive language on social media. Em *4th Italian Conference on Computational Linguistics (CLiC-it)*, 260–264. [↗](#)
- Phanontip, Aekachai, Thaiyathorn Sueb-in & Sirion Vittayakorn. 2021. Cyberbullying detection on Tweets. Em *18th International Conference on Electrical Engineering/Electronics, Computer, Telecommunications and Information Technology (ECTI-CON)*, 295–298. [doi](#) 10.1109/ECTI-CON51831.2021.9454848
- Poletto, Fabio, Valerio Basile, Manuela Sanguinetti, Cristina Bosco & Viviana Patti. 2020. Resources and benchmark corpora for hate speech detection: a systematic review. *Language Resources and Evaluation* 55. 477–523. [doi](#) 10.1007/s10579-020-09502-8
- Ribeiro, Marco, Sameer Singh & Carlos Guestrin. 2016. “why should I trust you?”: Explaining the predictions of any classifier. Em *Conference of the North American Chapter of the Association for Computational Linguistics (NAACL)*, 97–101. [doi](#) 10.18653/v1/N16-3020
- Rossi, Rafael Geraldelli. 2015. *Classificação automática de textos por meio de aprendizado de máquina baseado em redes*: Instituto de Ciências Matemáticas e de Computação, Universidade de São Paulo. Tese de Doutorado. [↗](#)
- Safikhani, Parisa & David Broneske. 2025. AutoML meets hugging face: Domain-aware pretrained model selection for text classification. Em *Conference of the Nations of the Americas Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL)*, 466–473. [doi](#) 10.18653/v1/2025.naacl-srw.45
- Saifullah, Shoffan, Rafał Dreżewski, Felix Andika Dwiyanto, Agus Sasmito Aribowo, Yuli Fauziah & Nur Heri Cahyana. 2024. Automated text annotation using a semi-supervised approach with meta vectorizer and machine learning algorithms for hate speech detection. *Applied Sciences* 14(3). [doi](#) 10.3390/app14031078
- Salles, Isadora, Francielle Vargas & Fabrício Benevenuto. 2025. HateBRXplain: A benchmark dataset with human-annotated rationales for explainable hate speech detection in Brazilian Portuguese. Em *31st International Conference on Computational Linguistics (COLING)*, 6659–6669. [↗](#)
- Santos, Raquel Bento, Bernardo Cunha Matos, Paula Carvalho, Fernando Batista & Ricardo Ribeiro. 2022. Semi-supervised annotation of portuguese hate speech across social media domains. Em *11th Symposium on Languages, Applications and Technologies (SLATE)*, 11:1–11:14. [doi](#) 10.4230/OASICS.SLATE.2022.11
- Schmidt, Anna & Michael Wiegand. 2017. A survey on hate speech detection using natural language processing. Em *5th International Workshop on Natural Language Processing for Social Media*, 1–10. [doi](#) 10.18653/v1/W17-1101
- da Silva Oliveira, Amanda, Thiago de Carvalho Cecote, João Paulo Reis Alvarenga, Vander Luis de Souza Freitas & Eduardo José da Silva Luz. 2024. Toxic speech detection in Portuguese: A comparative study of large language models. Em *15th International Conference on Computational Processing of Portuguese (PROPOR)*, 108–116. [↗](#)
- Smith, Noah A. 2020. Contextual word representations: putting words into computers. *Communications of the ACM* 63(6). 66–74. [doi](#) 10.1145/3347145

- Soto, Claver, Gustavo Nunes & José Gomes. 2019. Avaliação de técnicas de word embedding na tarefa de detecção de discurso de ódio. Em *Anais do XVI Encontro Nacional de Inteligência Artificial e Computacional (ENIAC)*, 1020–1031. doi [10.5753/eniac.2019.9354](https://doi.org/10.5753/eniac.2019.9354)
- Souza, Fábio, Rodrigo Nogueira & Roberto Lotufo. 2020. BERTimbau: Pretrained BERT models for Brazilian Portuguese. Em *Intelligent Systems*, 403–417. doi [10.1007/978-3-030-61377-8_28](https://doi.org/10.1007/978-3-030-61377-8_28)
- Suryawanshi, Shardul, Mihael Arcan & Paul Buitelaar. 2020. NUIG at SemEval-2020 task 12: Pseudo labelling for offensive content classification. Em *14th Workshop on Semantic Evaluation*, 1598–1604. doi [10.18653/v1/2020.semeval-1.208](https://doi.org/10.18653/v1/2020.semeval-1.208)
- Szcz Aleksander, esny, Maciej Markiewicz, Łukasz Radliński & Przemysław Kazienko. 2025. Leveraging positional bias of LLM in-context learning with class-few-shot and MajMin alternating ordering. Em *Computational Science – ICCS 2025*, 54–62
- Tan, Zhen, Dawei Li, Song Wang, Alimohammad Beigi, Bohan Jiang, Amrita Bhattacharjee, Mansoor Karami, Jundong Li, Lu Cheng & Huan Liu. 2024. Large language models for data annotation and synthesis: A survey. Em *Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 930–957. doi [10.18653/v1/2024.emnlp-main.54](https://doi.org/10.18653/v1/2024.emnlp-main.54)
- Trajano, Douglas, Rafael H Bordini & Renata Vieira. 2024. OLID-BR: offensive language identification dataset for Brazilian Portuguese. *Language Resources and Evaluation* 58. 1263–1289. doi [10.1007/s10579-023-09657-0](https://doi.org/10.1007/s10579-023-09657-0)
- Vargas, Francielle, Isabelle Carvalho, Fabiana Rodrigues de Góes, Thiago Pardo & Fabrício Benevenuto. 2022. HateBR: A large expert annotated corpus of Brazilian Instagram comments for offensive language and hate speech detection. Em *13th Language Resources and Evaluation Conference (LREC)*, 7174–7183. [↗](#)
- Vargas, Francielle, Fabiana Goés, Isabelle Carvalho, Fabrício Benevenuto & Thiago A S Pardo. 2021. Contextual-lexicon approach for abusive language detection. Em *International Conference on Recent Advances in Natural Language Processing (RANLP)*, 1438–1447. [↗](#)
- Wang, Shuohang, Yang Liu, Yichong Xu, Chengguang Zhu & Michael Zeng. 2021. Want to reduce labeling cost? GPT-3 can help. Em *Findings of the Association for Computational Linguistics: EMNLP*, 4195–4205. doi [10.18653/v1/2021.findings-emnlp.354](https://doi.org/10.18653/v1/2021.findings-emnlp.354)
- Waseem, Zeerak & Dirk Hovy. 2016. Hateful symbols or hateful people? predictive features for hate speech detection on Twitter. Em *NAACL Student Research Workshop*, 88–93. doi [10.18653/v1/N16-2013](https://doi.org/10.18653/v1/N16-2013)
- Wolf, Thomas, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Remi Louf, Morgan Funtowicz, Joe Davison, Sam Shleifer, Patrick von Platen, Clara Ma, Yacine Jernite, Julien Plu, Canwen Xu, Teven Le Scao, Sylvain Gugger, Mariama Drame, Quentin Lhoest & Alexander Rush. 2020. Transformers: State-of-the-art natural language processing. Em *Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 38–45. doi [10.18653/v1/2020.emnlp-demos.6](https://doi.org/10.18653/v1/2020.emnlp-demos.6)
- X. 2026. The x rules. Acesso em: fev. 2026. [↗](#)
- Yang, An, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan Li, Dayiheng Liu, Fei Huang, Haoran Wei, Huan Lin, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang, Jiaxi Yang, Jingren Zhou, Junyang Lin, Kai Dang, Keming Lu, Keqin Bao, Kexin Yang, Le Yu, Mei Li, Mingfeng Xue, Pei Zhang, Qin Zhu, Rui Men, Runji Lin, Tianhao Li, Tianyi Tang, Tingyu Xia, Xingzhang Ren, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yu Wan, Yuqiong Liu, Zeyu Cui, Zhenru Zhang & Zihan Qiu. 2025. Qwen2.5 technical report. ArXiv [cs.CL]. doi [10.48550/arXiv.2412.15115](https://doi.org/10.48550/arXiv.2412.15115)
- Youtube. 2026. Hate speech policy. Acesso em: fev. 2026. [↗](#)
- Zampieri, Marcos, Shervin Malmasi, Preslav Nakov, Sara Rosenthal, Noura Farra & Ritesh Kumar. 2019. Predicting the type and target of offensive posts in social media. Em *Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL)*, 1415–1420. doi [10.18653/v1/N19-1144](https://doi.org/10.18653/v1/N19-1144)
- Zhang, Chuxu, Dongjin Song, Chao Huang, Ananthram Swami & Nitesh V. Chawla. 2019. Heterogeneous graph neural network. Em *25th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 793–803. doi [10.1145/3292500.3330961](https://doi.org/10.1145/3292500.3330961)
- Zhou, Dengyong, Olivier Bousquet, Thomas Navin Lal, Jason Weston & Bernhard Schölkopf.

2004. Learning with local and global consistency. Em *Advances in Neural Information Processing Systems (NeurIPS)*, 321–328. [↗](#)

Zhou, Zhi-Hua. 2012. *Ensemble methods: Foundations and algorithms*. CRC Press

Zhu, Xiaojin, Zoubin Ghahramani & John Lafferty. 2003. Semi-supervised learning using gaussian fields and harmonic functions. Em *20th International Conference on International Conference on Machine Learning*, 912–919. [↗](#)

Apêndices

A. Resultados detalhados do método semi-supervisionado

Corpus	Regularização	% Dados	Classificador	F-measure
ToLD-BR	GFHF	30%	MLP	0.64
	LGC	10%	SVM	0.72
HateBR	GFHF	30%	MLP	0.36
	LGC	30%	MLP	0.37
HLPDSD	GFHF	30%	MLP	0.26
	LGC	10%	SVM	0.41

Tabela 11: Melhores desempenhos do método semi-supervisionado por *corpus* e algoritmo de regularização, em termos de *F-measure*. Para cada caso, apresenta-se a configuração que obteve o melhor resultado.

B. Resultados detalhados dos classificadores baseados em LLMs

C. Lista de Termos Identitários

A lista de termos utilizada para identificar instâncias com conteúdo identitário foi construída por tradução e adaptação para o português brasileiro da lista proposta por Dixon et al. (2018); Kennedy et al. (2020), com adição de termos relacionados ao contexto sociopolítico brasileiro, incluindo jargões associados ao debate político nacional do período de coleta dos *corpora*. Os termos estão organizados por categoria para facilitar a leitura, mas foram tratados de forma agregada nos experimentos, isto é, uma instância é considerada identitária se contém ao menos um termo de qualquer categoria.

Política petralha, petista, mortadela, bolsomion, bozo, bolsonarista, esquerdopata, fascista, nazista, comunista, gado, lulista, esquerda, direita, tucano

Corpus	Método	Modelo	Original		Filtrado	
			F	κ	F	κ
ToLD-BR	zero-shot	Qwen2.5-7B	0,73	0,41	0,74	0,44
		Granite3.3-8B	0,73	0,47	0,80	0,55
	few-shot	Qwen2.5-7B	0,74	0,44	0,79	0,49
		Granite3.3-8B	0,74	0,48	0,81	0,56
	RAG	Qwen2.5-7B	0,74	0,45	0,78	0,45
		Granite3.3-8B	0,73	0,47	0,80	0,50
HateBR	zero-shot	Qwen2.5-7B	0,86	0,71	0,87	0,73
		Granite3.3-8B	0,75	0,56	0,82	0,64
	few-shot	Qwen2.5-7B	0,87	0,73	0,89	0,78
		Granite3.3-8B	0,80	0,64	0,86	0,74
	RAG	Qwen2.5-7B	0,85	0,70	0,87	0,73
		Granite3.3-8B	0,83	0,67	0,88	0,75
HLPDSD	zero-shot	Qwen2.5-7B	0,62	0,38	0,82	0,61
		Granite3.3-8B	0,58	0,40	0,70	0,40
	few-shot	Qwen2.5-7B	0,63	0,41	0,83	0,64
		Granite3.3-8B	0,59	0,40	0,80	0,62
	RAG	Qwen2.5-7B	0,61	0,39	0,82	0,63
		Granite3.3-8B	0,60	0,39	0,81	0,61

Tabela 12: Desempenho dos classificadores baseados em LLMs com os modelos Qwen2.5-7B e Granite3.3-8B nos *corpora* avaliados, em termos de *F-measure* (F) e coeficiente Kappa de Cohen (κ), nos cenários dos *corpora* originais e filtrados.

Gênero e Sexualidade mulher, lésbica, gay, bissexual, transgênero, trans, queer, lgbt, homossexual, bicha, viado, sapatão, travesti, fêmea, noiva, esposa, mãe

Raça e Etnia negro, preto, pardo, branco, africano, índio, afro-americano, latino, asiático, indígena, cigano, macaco, senzala, escravo

Religião muçulmano, judeu, islã, islâmico, alá, judaico, cristão, católico, evangélico, clero, maomé, religião

Deficiência e Idade cego, surdo, mudo, paralisado, deficiente, autista, retardado, velho, idoso, jovem, adolescente

D. Métricas complementares de desempenho

Corpus	Método	Configuração	Acurácia		Precisão		Cobertura		F-measure		Macro F-measure	
			Orig.	Filt.	Orig.	Filt.	Orig.	Filt.	Orig.	Filt.	Orig.	Filt.
ToLD-BR	BERT	BERTimbau FT	0.72	0.77	0.64	0.71	0.84	0.90	0.72	0.80	0.72	0.77
	<i>zero-shot</i>	Granite3.3-8B	0.73	0.73	0.66	0.66	0.83	0.83	0.73	0.73	0.73	0.73
	semi	LGC_10%_SVM	0.73	0.78	0.66	0.76	0.81	0.81	0.72	0.79	0.73	0.78
	<i>few-shot</i>	Granite3.3-8B	0.74	0.78	0.66	0.72	0.83	0.93	0.74	0.81	0.74	0.78
	RAG	Granite3.3-8B / Qwen	0.73	0.75	0.64	0.68	0.85	0.96	0.74	0.80	0.73	0.74
	Comitê	-	0.75	0.79	0.66	0.72	0.89	0.95	0.76	0.82	0.75	0.78
HateBR	BERT	BERTimbau FT	0.83	0.87	0.86	0.90	0.79	0.83	0.83	0.86	0.83	0.87
	<i>zero-shot</i>	Qwen2.5-7B	0.86	0.87	0.84	0.83	0.88	0.94	0.86	0.88	0.85	0.87
	semi	LGC_30%_MLP	0.61	0.64	0.97	0.96	0.23	0.28	0.37	0.44	0.54	0.59
	<i>few-shot</i>	Qwen2.5-7B	0.87	0.89	0.86	0.87	0.87	0.92	0.87	0.89	0.87	0.89
	RAG	Granite3.3-8B / Qwen	0.83	0.87	0.86	0.84	0.80	0.92	0.83	0.88	0.84	0.87
	Comitê	-	0.86	0.90	0.88	0.91	0.82	0.88	0.85	0.89	0.86	0.90
HLPBSD	BERT	BERTimbau FT	0.73	0.77	0.58	0.89	0.55	0.62	0.56	0.73	0.68	0.76
	<i>zero-shot</i>	Qwen2.5-7B	0.70	0.81	0.51	0.77	0.79	0.87	0.62	0.82	0.68	0.80
	semi	LGC_10%_SVM	0.53	0.61	0.69	0.95	0.47	0.24	0.41	0.38	0.42	0.55
	<i>few-shot</i>	Qwen2.5-7B	0.72	0.82	0.54	0.81	0.76	0.85	0.63	0.83	0.70	0.82
	RAG	Qwen2.5-7B	0.72	0.81	0.54	0.78	0.69	0.87	0.61	0.82	0.70	0.81
	Comitê	-	0.74	0.83	0.58	0.83	0.68	0.83	0.62	0.83	0.71	0.83

Tabela 13: Comparação do desempenho dos métodos nos corpora ToLD-BR, HateBR e HLPBSD nos cenários original (*Orig.*) e filtrado (*Filt.*), considerando acurácia, precisão, cobertura, *F-measure* e *Macro F-measure*.

<http://www.linguamatica.com/>

Artigos de Investigação

A Transcrição Fonética como Ferramenta Pedagógica no Ensino de Português para Hispanofalantes: Um Estudo De Caso

Keila Coelho

Norma, ús i interferència: biaixos lingüístics en els models de llenguatge en català

Mireia Almena Rodríguez & Thomas Brochhagen

Um Comitê de Classificadores para Anotação Automática de Linguagem Tóxica em Português sob Escassez de Dados

F. A. R. Neto, R. T. Anchiêta, R. S. Moura, P. A. S. Neto & A. M. Santana