

Diálogos Tóxicos: Gatilhos e Padrões de Interação no Reddit Brasileiro

Giovana Piorino Vieira de Carvalho
Departamento de Ciência da Computação
Universidade Federal de Minas Gerais
Belo Horizonte, Brasil
giovana.p@dcc.ufmg.br

Ana Paula Couto da Silva
Departamento de Ciência da Computação
Universidade Federal de Minas Gerais
Belo Horizonte, Brasil
ana.coutosilva@dcc.ufmg.br

Abstract—The increasing volume in user generated content challenges the understanding of online toxicity’s structural dynamics. Despite advances in detection models, few studies analyze the discursive architecture distinguishing toxic from non-toxic interactions or the role of triggers, especially in low-resource languages like Brazilian Portuguese. We address this by investigating structural and linguistic patterns in Reddit discussion trees, focusing on how triggers escalate conflicts. We fine-tuned BERTabaporu, optimized via data augmentation, to classify comments and reconstruct discussion trees. A multidimensional analysis integrating vector similarity, morphosyntactic tagging, and Named Entity Recognition mapped the conversations’ evolution. Results show that toxic trees exhibit greater depth and engagement, with initial semantic cohesion that degrades over time. Unlike non-toxic discussions, which prioritize social bonding, toxic threads are driven by political polarization and argumentative syntax. This study elucidates toxicity scaling, supporting the development of predictive mechanisms for early identification of discursive escalation.

I. INTRODUÇÃO

Nos últimos anos, as plataformas de mídia social online se tornaram ferramentas populares para interações globais, principalmente devido a natureza intrinsecamente interativa que promove um senso de pertencimento e conexão emocional [1]. Estas plataformas passaram a ser um espaço para troca de informações e busca de apoio emocional para problemas do dia a dia.

Infelizmente, nem sempre estas interações são positivas. Estudos mostram que aproximadamente 5-10% das interações online são classificadas como tóxicas na maioria das plataformas moderadas, podendo atingir até 19% em plataformas não moderadas [2]. Estas interações tóxicas são caracterizadas por linguagem ofensiva, desrespeitosa ou emocionalmente prejudicial, o que pode afetar o engajamento dos usuários, diminuir a confiança dentro das comunidades e apresentar riscos à saúde mental de indivíduos mais vulneráveis [3]. Como exemplo, a pesquisa realizada com 29.593 estudantes universitários em Arif et al. [4] fornece evidências da prevalência e do impacto desse problema na sociedade contemporânea, mostrando que a depressão e a ansiedade estão fortemente associadas a experiências de *cyberbullying*, uma forma de comportamento tóxico online.

Entre as diversas plataformas, a rede social *Reddit*¹ se destaca por suas dinâmicas de interação pautada em tópicos de discussão, bem como sua moderação voluntária. Por um lado, essa estrutura facilita discussões abertas e o compartilhamento, de forma anônima, de experiências pessoais [5]. Porém, apresenta desafios em termos de moderação e controle de conteúdo tóxico: estudos indicam que 16% dos usuários publicam postagens e comentários tóxicos, com 81% destes variando sua conduta entre comunidades para se adequarem às suas normas [6].

O comportamento tóxico de um ou mais usuários pode emergir em qualquer nível dentro de uma árvore de discussão (ou *thread*), mesmo em interações inicialmente positivas, fenômeno analisado em diferentes estudos da literatura [7, 8, 9]. Esses trabalhos conduzem o conceito de gatilhos: comentários *não tóxicos* que provocam respostas *tóxicas* na estrutura das discussões. Identificar esses elementos é fundamental para a moderação no *Reddit*, permitindo aos moderadores adotar uma postura proativa frente à potenciais cenários de toxicidade.

Neste contexto, o objetivo principal deste trabalho é identificar e analisar árvores de discussão com e sem gatilho em comunidades do *Reddit* (em português brasileiro). Combinando propriedades linguísticas e métricas estruturais, buscamos identificar padrões que distingam elementos desencadeadores de toxicidade de dinâmicas conversacionais saudáveis, considerando o contexto de comunidades brasileiras na rede social. Nosso trabalho expande os resultados apresentados em [10], buscando superar suas principais limitações: o alto custo atrelado ao uso de modelo proprietário e o tamanho reduzido da amostra de árvores de discussão. Para isso, propomos uma abordagem baseada em um modelo de código aberto, ampliamos consideravelmente o volume de árvores analisadas e incorporamos novas ferramentas de processamento linguístico, proporcionando uma análise mais robusta das discussões e estrutura de comentários gatilho.

Nossas principais contribuições são as seguintes: (i) disponibilização de um conjunto de dados contendo 883.953 árvores de discussão do *Reddit* brasileiro; (ii) modelo de classificação de toxicidade para português brasileiro; e (iii) análises topológicas e linguísticas da toxicidade no contexto

¹<https://www.reddit.com/>

brasileiro. Os resultados são importantes para a proposta de mecanismos de moderação automática, como ferramenta complementar para mitigar comportamentos nocivos em plataformas sociais online.

II. TRABALHOS RELACIONADOS

Estudos sobre toxicidade priorizam a classificação de comentários e a caracterização linguística, com menor ênfase no impacto das discussões e em caracterização de gatilhos. Enfatizando o contexto do *Reddit*, observa-se predominância de pesquisas em inglês [11, 12], sendo que modelos voltados à detecção e caracterização de toxicidade, bem como análises de impacto de comentários gatilho em discussões em português brasileiro [13, 14, 15] ainda apresentam menor exploração.

Sob a ótica da dinâmica e métricas de toxicidade, estudos recentes focam na estrutura das conversas, no contexto do *Reddit*. Em Shankaran and Sharma [16], o modelo *ToxiGen* foi empregado para analisar o espectro de toxicidade em textos de língua inglesa, concluindo que a hostilidade de um comentário é predominantemente determinada pelo seu antecessor imediato, corroborando achados de Maleki et al. [3]. Neste último, utilizando o modelo *Detoxify*, observou-se que discussões tóxicas tendem a ser mais longas, profundas e frequentemente terminam em toxicidade.

No contexto de detecção de gatilhos, definidos como comentários não tóxicos que suscitam respostas tóxicas, Almerexhi et al. [8] utilizaram um classificador LSTM com representações GloVe, obtendo F1-Score de 0.83, também no contexto de discussões do *Reddit* em língua inglesa. Os autores concluem que mudanças abruptas de tópico e termos políticos atuam como fortes indicadores locais de gatilhos. Avançando nessa abordagem, os autores em Almerexhi et al. [9] propõem o modelo PROVOKE (Bi-LSTM), que integra *embeddings* textuais a atributos contextuais, como variação de sentimento, tópico e profundidade na árvore. Alcançando um F1-Score de 0.78, o estudo destaca que, além de termos ofensivos genéricos, o vocabulário específico de cada comunidade desempenha papel relevante na caracterização dos gatilhos.

Embora estudos anteriores tenham investigado como comentários inofensivos podem desencadear toxicidade em discussões online, essas pesquisas concentram-se predominantemente em contextos de língua inglesa. Nosso trabalho preenche essa lacuna ao analisar aspectos linguísticos que sinalizam potenciais gatilhos de toxicidade em interações no *Reddit* brasileiro, um contexto ainda pouco explorado na literatura.

III. METODOLOGIA

Esta seção descreve a metodologia utilizada neste trabalho.

A. Conjunto de Dados

O conjunto de dados analisado neste trabalho foi compartilhado publicamente em Lima et al. [17] e é formado por postagens e comentários feitos entre janeiro e dezembro de 2022 nos dez subreddits brasileiros mais populares neste período: *brasil*, *desabafos*, *futebol*, *saopaulo*, *eu_nvr*, *botecodoredit*,

conversas, *investimentos*, *tiodopave*, *brasilivre*. A sua coleta foi realizada através da API Pushshift.²

Resumidamente, os dados analisados incluem comentários apenas em português, excluindo comentários de comunidades que permitem discussões em múltiplos idiomas. Para o resultado final, comentários categorizados como deletados ou removidos foram filtrados, juntamente com comentários contendo apenas emojis ou símbolos, URLs, caracteres não alfanuméricos, e reações de texto apenas de risadas.³ Finalmente, foram excluídos comentários gerados por contas de moderação automática e bots que detectamos em nossos dados. Assim, nosso corpus, após a aplicação dos filtros, possui 6.589.541 milhões de comentários e aproximadamente 390.000 postagens.

B. Anotação Manual

Para a amostra rotulada manualmente, inicialmente foram selecionados 2.500 comentários de forma estratificada. Estes comentários foram divididos em cinco lotes de 500 comentários cada. Doze estudantes de graduação e pós-graduação em Ciência da Computação e Letras foram divididos em 4 grupos de anotação, sendo que o grupo que teve maior concordância entre os membros anotou dois lotes. Cada comentário foi classificado em uma das categorias: *Tóxico*, *Não Tóxico*, *Não sei*, ou *Informação insuficiente*, seguindo as definições estabelecidas pela API Perspective⁴, que foram incluídas no conjunto de instruções fornecidas aos participantes do processo de anotação manual [17].

Como esperado, a porcentagem de comentários *Tóxicos* foi muito inferior à de *Não Tóxicos*. Para aumentarmos a quantidade de comentários anotados manualmente para esse rótulo, realizamos uma pré-anotação automática com a Perspective API⁵, selecionando comentários considerados *Tóxicos* pelo modelo, e realizamos uma validação com um dos grupos de anotadores que participaram na primeira etapa de rotulação acerca do rótulo atribuído pela Perspective.

A partir dos resultados das rodadas de anotação, consideramos apenas os comentários cuja concordância entre pelo menos dois anotadores foi atribuída aos rótulos *Não tóxico* ou *Tóxico* (voto majoritário), resultando em um total de 2.961 comentários de interesse. Deles, 85,21% foram classificados como *Não tóxico*, e 14,79% como *Tóxico*. Tais comentários rotulados foram utilizados para o treinamento e validação do modelo de classificação de toxicidade proposto.

C. Modelo de Toxicidade para Português do Brasil

Para a rotulação automática de comentários como *Não tóxico* ou *Tóxico*, realizamos o ajuste fino do BERTabaporu, um modelo de linguagem para o português brasileiro construído sobre a arquitetura BERT e pré-treinado em aproximadamente 238 milhões de tweets [18]. A escolha deste

²<https://github.com/pushshift/api>.

³Em português, textos de risadas são representados pela sequência de caracteres kkkk

⁴https://developers.perspectiveapi.com/s/about-the-api-attributes-and-languages?language=en_US

⁵<https://perspectiveapi.com/>

modelo foi motivada pela similaridade linguística entre seu domínio de pré-treinamento e nosso corpus alvo, sendo ambos textos informais de mídias sociais. Além disso, a ausência de custos financeiros torna essa solução mais viável em comparação aos Grandes Modelos de Linguagem (LLMs), especialmente diante da extensa base de dados a ser rotulada [10].

Para adaptar o modelo para nossa tarefa de classificação binária de toxicidade, utilizamos a classe *AutoModelForSequenceClassification* da biblioteca Transformers do Hugging Face [19].

Antes do treinamento, os comentários passaram por um pré-processamento de texto padrão para remover ruídos típicos de conteúdo de mídia social. Todo o texto foi convertido para minúsculas, e URLs, menções a usuários, hashtags, emojis, pontuação, números e stopwords foram removidos. *Stopwords* foram definidas de acordo com a lista de stopwords do português fornecida pelo corpus NLTK [20], garantindo consistência com as práticas de pré-processamento comumente adotadas na classificação de textos de mídia social [21, 22].

Visando mitigar o desbalanceamento das classes, adotamos a estratégia de aumento de dados anotados manualmente por humanos utilizando grandes modelos de linguagem (LLMs). Selecionamos 5.000 novos comentários do conjunto de dados, que foram classificados por 3 LLMs: Sabiá-2 (especializado em português) [23], Llama-2[24] e Mistral [25]. Somente os comentários com voto majoritário entre os modelos e aqueles que não foram previamente rotulados por humanos foram incluídos no conjunto de anotação automática (sintética). Assim, foram adicionados 3.385 comentários aos dados descritos na Seção 3.2. O prompt utilizado para a tarefa foi proposto em [17].

Os dados sintéticos foram incorporados exclusivamente ao conjunto de treinamento. O conjunto de teste é composto somente por dados anotados por humanos, garantindo uma avaliação livre de possíveis vieses introduzidos pelos grandes modelos de linguagem. A Tabela I apresenta os dados usados para definição do modelo. A inclusão dos dados sintéticos⁶ quase dobrou os exemplos para treinamento, enriquecendo a classe de comentários *Tóxicos* cuja a acurácia de classificação é importante para a identificação da presença de gatilhos nas discussões analisadas. O modelo final foi disponibilizado publicamente⁷.

Os hiperparâmetros de taxa de aprendizado e tamanho do lote foram definidos a partir de uma busca em grade, conforme mostrado na Tabela II, considerando a combinação dos valores $\{1 \times 10^{-5}, 2 \times 10^{-5}, 3 \times 10^{-5}, 5 \times 10^{-5}\}$ e $\{16, 32, 64\}$, respectivamente. O modelo foi treinado usando validação cruzada com 5 *folds*. A configuração que maximizou o desempenho foi a taxa de aprendizado de 3×10^{-5} com tamanho de lote de 16, utilizando o otimizador AdamW e *Early Stopping* de 3 épocas.

⁶Foram filtrados comentários com rótulos *Não sei dizer*, *Falta Informação* e comentários duplicados

⁷<https://huggingface.co/giovanaPcarvalho/bertabaporu-toxicity-reddit-ptbr>

TABLE I
DISTRIBUIÇÃO DE COMENTÁRIOS ROTULADOS POR FONTE E CLASSE.

Conjunto	Fonte da Anotação	Tóxicos	Não tóxicos	Total
Treino	Humanos	350	2.018	2.368
	Sintética (LLMs)	535	2.850	3.385
	<i>Subtotal Treino</i>	<i>885</i>	<i>4.868</i>	<i>5.753</i>
Teste	Humanos	88	505	593

Taxa de aprendizado ($\times 10^{-5}$)	Tamanho de lote	F1-score	Recall	Precisão
1	16	0.8877	0.8797	0.9073
1	32	0.8949	0.8877	0.9103
1	64	0.8881	0.8816	0.9135
2	16	0.9022	0.8962	0.9150
2	32	0.8958	0.8875	0.9157
2	64	0.8944	0.8870	0.9113
3	16	0.9133	0.9094	0.9212
3	32	0.8906	0.8820	0.9114
3	64	0.8927	0.8842	0.9122
5	16	0.8825	0.8719	0.9096
5	32	0.8801	0.8693	0.9096
5	64	0.8887	0.8797	0.9104

TABLE II
DESEMPENHO MÉDIO DE TODAS AS COMBINAÇÕES DE TAXA DE APRENDIZADO E TAMANHO DE LOTE DURANTE A VALIDAÇÃO CRUZADA ESTRATIFICADA COM 5 FOLDS.

D. Construção das Árvores de Discussão

Cada comentário coletado contém identificadores únicos, como o *id*, o *link_id* (que indica a postagem principal à qual pertence) e o *parent_id* (que aponta para o comentário imediatamente anterior na hierarquia). A partir dessas relações, é possível reconstruir a estrutura das conversas, conectando cada resposta ao seu respectivo comentário pai, gerando as sequências de discussões completas. Assim, consideramos como árvore de discussão toda a troca de comentários que se inicia a partir de uma postagem, seguida por um comentário resposta à ela, que é um comentário raiz. Respostas diretas e subsequentes ao comentário raiz são estruturadas e compõem uma árvore de discussão.

As árvores de discussão são divididas em dois grupos: *Não Tóxicas* e *Tóxicas*. Árvores *Não Tóxicas* são aquelas em que todos os comentários possuem rótulo *Não Tóxico*. Árvores *Tóxicas* possuem comentários *Tóxicos* e *Não Tóxicos*, mas possuem a seguinte característica: o comentário raiz deve ter rótulo *Não Tóxico*. Esta condição garante que a toxicidade na discussão sempre surja como uma resposta a um conteúdo *Não Tóxico*, definindo assim um gatilho de discussão. Importante notar que o gatilho pode estar em qualquer posição da árvore de discussão, uma vez que o comentário *Não Tóxico* deverá ser seguido por um comentário *Tóxico*, não sendo necessariamente a raiz da árvore. Árvores de discussão que não atendem a essa condição (como as 100% *Tóxicas* ou as iniciadas por uma raiz *Tóxica*) não são consideradas em nossas análises.

Nosso objetivo é identificar traços implícitos de linguagem que possam contribuir para mudanças de rumos nas discussões

e identificar proativamente quando uma discussão possui o potencial de se tornar tóxica. Além disso, a partir dessa estruturação, comparamos estatísticas de profundidade entre os grupos de árvores de discussão e, em árvores *Tóxicas*, avaliar proporções de respostas tóxicas após um gatilho.

E. Caracterização Linguística

Para encontrar diferenças e similaridades entre árvores de discussão com gatilho (*Tóxicas*) e sem gatilho (*Não Tóxicas*), realizamos o conjunto de análises linguísticas descritas a seguir.

Similaridade Semântica. Para comparar a semântica dos comentários analisados, construímos um espaço semântico de alta dimensão S^n . Cada comentário analisado é incorporado neste espaço através de sua representação vetorial. A codificação no espaço semântico é feita através do modelo de *embedding paraphrase-multilingual-MiniLM-L12-v2*, um SBERT (Sentence-BERT) multilíngue, presente na biblioteca *Hugging Face*.⁸ A escolha do modelo se deve ao fato dele ser treinado especificamente em tarefas de similaridade com suporte à língua portuguesa [26].

Apesar de especificado para agrupar conteúdos com significado similar em regiões próximas, o espaço semântico S^n também captura, de forma indireta, similaridade de assuntos tratados, estruturas discursivas (retórica), padrões de argumentação, estilo e contexto. Ou seja, comentários similares, no sentido mais amplo da definição, estarão próximos no espaço vetorial construído. Assim, para quantificar o grau de proximidade entre pares de comentários adjacentes (um comentário de uma árvore de discussão e sua resposta imediata), calculamos a métrica de similaridade de cosseno.

Nosso interesse é observar quando a discussão em andamento muda significativamente, ou seja, possuem baixa similaridade dentro do espaço semântico construído. Para tal, definimos um limiar para determinar mudanças na similaridade entre pares de comentários, podendo ser um *proxy* para mudança de foco e estilo da discussão em análise, seguindo [27]. Este limiar é calculado usando a média e o desvio-padrão de todos os valores de similaridade (comentários par a par) das árvores analisadas e foi definido como 1 desvio padrão abaixo da média para considerar quedas de similaridade estatisticamente significativas [27]. Pares de comentários com similaridade abaixo desse limiar foram considerados uma tupla de mudança abrupta de discurso.

Reconhecimento de Entidades Nomeadas (REN). Nessa análise, aplicamos o componente REN, modelo apresentado na biblioteca spaCy, adaptando-o aos nossos dados por meio do componente de POS tagging junto ao corpus WikiNER [28]. Essa abordagem classifica as entidades em quatro categorias: Pessoa (PER), Localização (LOC), Organização (ORG) e Diversos (MISC) e foi utilizada para identificar os principais atores citados nas interações, considerando que os dados

abrangem um ano eleitoral e de grandes questões geopolíticas.

Etiquetagem de classe de palavra (Pos Tagging). Para identificar as classes gramaticais predominantes em cada tipo de árvore de discussão, aplicamos POS tagging [29] com o modelo⁹ disponível na biblioteca spaCy [30], baseado em uma treebank [31] e anotada segundo o padrão Universal Dependencies [32]. Entre as possíveis funções do modelo, utilizamos especificamente o componente morfologizador, responsável pela atribuição de categorias gramaticais. Adicionalmente, expandimos a análise para a avaliação de n-gramas sintáticos (bigramas e trigramas) a partir das sequências de etiquetas. Essa abordagem, baseada em exemplos de estudos relacionados à n-gramas [33], permite capturar padrões estruturais gramaticais independentemente do conteúdo semântico ou tópico da discussão.

Classificação de Tipo de Discurso de Ódio. A classificação das categorias de discurso de ódio foi realizada utilizando o modelo *pysentimiento/bertabaporu-pt-hate-speech*, treinado para detectar diferentes tipos de discurso em Português [34, 35] e disponível para uso por meio da biblioteca *Hugging Face*.¹⁰ Fundamentada na arquitetura RoBERTa, a abordagem categoriza cada comentário em cinco classes principais: Sexismo (*Sexism*), Homofobia (*Homophobia*), Racismo (*Racism*), Ideologia (*Ideology*) calculando-se a probabilidade individual de enquadramento em cada rótulo, para cada comentário. O modelo foi aplicado exclusivamente aos comentários pertencentes às árvores *Tóxicas*, de forma a categorizar as tipologias de discurso nocivo. Com base em [34], os comentários foram rotulados com as categorias nas quais o modelo atribui probabilidade maior ou igual a 0,5 de presença de conteúdo associado.

IV. RESULTADOS

A seguir apresentamos os resultados do trabalho.

A. Anotação Automática de Toxicidade

A Tabela III apresenta o desempenho do modelo de toxicidade descrito na Seção III-C e aplicado no conjunto de teste composto por 593 comentários anotados unicamente por humanos. Os resultados mostram que o modelo ajustado para a tarefa de toxicidade tem um bom desempenho ao classificar o conteúdo tóxico, principalmente quando comparado com outros modelos abertos apresentados na literatura [21], alcançando um F1-Score Médio Macro de 0,80. Vale ressaltar que o nosso objetivo neste trabalho é ter um modelo aberto, escalável e com desempenho satisfatório, a ser aplicado no grande conjunto de comentários das árvores de discussão a serem analisadas.

Complementando as métricas quantitativas de desempenho, a seguir apresentamos dois exemplos de comentários onde o nosso modelo divergiu em relação à rotulação humana.

⁸<https://huggingface.co/sentence-transformers/paraphrase-multilingual-MiniLM-L12-v2>

⁹*pt_core_news_lg*

¹⁰<https://huggingface.co/pysentimiento/bertabaporu-pt-hate-speech>

Classe	Precisão	Recall	F1-Score
Não Tóxico (0)	0,93	0,95	0,94
Tóxico (1)	0,70	0,61	0,65
Acurácia			0,90
Média Macro	0,82	0,78	0,80
Média Ponderada	0,90	0,90	0,90

TABLE III

DESEMPENHO DO MODELO DE TOXICIDADE.

O primeiro exemplo é um *falso negativo*, onde o modelo falhou em detectar toxicidade explícita, apesar do uso de adjetivos ofensivos direcionados a uma figura pública: “*mano esse luiz inacio ta ficando velho senil e psicopata na moral hahahaha*”. A inclusão de marcadores discursivos informais e risos atenua superficialmente o tom agressivo, provavelmente influenciando o modelo a interpretar o enunciado como não hostil. Esse comportamento é consistente com limitações observadas em tarefas que exigem inferência pragmática: expressões de humor e marcações coloquiais podem suavizar a superfície linguística sem remover o conteúdo ofensivo [36, 37].

Por outro lado, o modelo apresentou casos de *falso positivos*, em que classificou indevidamente como tóxicos comentários que continham linguagem intensa, porém sem ataque interpessoal. Como exemplo, temos o comentário “*to bem f*-se, quero que o pau quebre*”, marcado pelo uso de palavras de xingamento, para expressar um possível sentimento de frustração. No entanto, modelos de toxicidade frequentemente associam palavras à categoria tóxica de forma correlacional, sem distinguir entre xingamentos direcionados e manifestações gerais de emoção, resultando neste tipo de falso positivo.

Por fim, o modelo proposto foi aplicado ao conjunto de dados composto por 2.486.593 comentários das árvores candidatas a pertencerem às classes *Tóxica* e *Não Tóxica*. O total de comentários rotulados automaticamente como *Tóxicos* foi de 271.235 (10,91%) e como *Não Tóxicos* foi de 2.215.358 (89,09%), o que é consistente com a prevalência esperada de toxicidade em grandes comunidades online [38].

B. Topologia das Árvores de Discussão

O modelo proposto foi utilizado para rotular os 2.486.593 comentários das árvores candidatas a pertencerem às classes *Tóxica* e *Não Tóxica*. Vale ressaltar que o critério para classificar uma árvore como *Tóxica* observa se o par subsequente de comentários tem um gatilho (comentário *Não Tóxico*) e uma resposta *Tóxica*, mesmo se o restante das respostas subsequentes forem rotuladas como *Tóxicas* ou *Não Tóxicas*. Após a validação dos critérios de classificação, o conjunto de árvores para análise consiste em 883.553 árvores de discussão, das quais 813.445 (92,07%) são *Não Tóxicas* e 70.508 (7,98%) são *Tóxicas*.

Como avaliado na Tabela IV, todas as métricas estatísticas de profundidade avaliadas (média, mediana, máximo e desvio-padrão) revelaram-se superiores em árvores *Tóxicas*.

Grupo de árvores	Média	Desvio Padrão	Mediana	Q3 (75%)	Máximo
Não Tóxicas	2.72	1.25	2.00	3.00	79
Tóxicas	3.85	2.39	3.00	5.00	84

TABLE IV

ESTATÍSTICAS DESCRITIVAS DA PROFUNDIDADE DAS ÁRVORES DE DISCUSSÃO *Tóxicas* E *Não Tóxicas*.

A elevação da mediana e do terceiro quartil sugere que gatilhos fomentam réplicas sucessivas, postergando o encerramento rápido observado em árvores *Não Tóxicas*. Além disso, o maior desvio-padrão (2,39 em árvores *Tóxicas* contra 1,25 em árvores *Não Tóxicas*) indica maior dispersão estrutural em discussões *Tóxicas*.

A análise estrutural da discussão *Tóxica* após a ocorrência de um gatilho mostra que o volume médio de comentários subsequentes ao gatilho é de 2,11, representando em média 53,29% do volume total de comentários da discussão. Isso sugere que o gatilho não é um evento final isolado, e ocorre tipicamente na metade de uma interação média. Além disso, o gatilho provoca um impacto significativo no andamento da conversa: as respostas subsequentes apresentam uma taxa média de toxicidade de 73,67%, revelando uma mudança expressiva no tom do diálogo após sua ocorrência. No entanto, essa escalada não se sustenta de forma contínua. Ao analisar sequências de respostas consecutivamente tóxicas após o gatilho, observamos um encadeamento médio de apenas 1,11, indicando que essas manifestações tendem a surgir de maneira pontual. Esse padrão sugere variações importantes no modo como os usuários reagem tanto ao gatilho quanto ao desenrolar da discussão.

C. Características linguísticas

Para a formulação das análises linguísticas, removemos ocorrências de árvores de discussão de profundidade dois, ou seja, apenas um comentário inicial com uma resposta, por não agregarem uma estrutura interessante em termos de evolução da discussão. Assim, após esse filtro, temos 321.174 árvores de discussão *Não Tóxicas* (compreendendo 1.230.816 comentários), e 47.860 árvores de discussão *Tóxicas* (com 225.939 comentários).

Similaridade Semântica. No nosso conjunto de dados, a média de similaridade dos 1.087.721 pares de comentários adjacentes foi igual a 0,395, sendo considerado como 0,21 o limiar de mudança de conteúdo entre interações sucessivas. Este valor foi utilizado para as análises descritas a seguir.

Considerando os conjuntos de árvores separadamente, como mostrado na Tabela V, para árvores de discussão *Tóxicas*, temos taxas de mudanças de similaridade menores (14,84%) em comparação às árvores *Não Tóxicas* (17,29%). Isso sugere que as discussões tóxicas tendem a fixar mais o tópico da conversa. Além disso, a análise específica do par *Gatilho* → *Resposta* reforça essa observação, em que 14,62% desses pares apresentam similaridade abaixo do limiar estabelecido, taxa ligeiramente menor que a do conjunto de árvores *Não Tóxicas*, indicando que o gatilho e sua resposta imediata apresentam

Faixa de Profundidade da Árvore	Árvores Não Tóxicas (% com Mudança)	Árvores Tóxicas (% com Mudança)
3-5	35,40	33,69
6-10	54,00	53,36
11-20	65,60	67,17
21-50	68,57	82,35
51+	100,00	100,00

TABLE V

PORCENTAGEM DE ÁRVORES DE DISCUSSÃO CONTENDO PELO MENOS UMA MUDANÇA DE TÓPICO, CATEGORIZADAS POR PROFUNDIDADE.

Pessoas (PER)		Organizações (ORG)		Locais (LOC)	
Entidade	%	Entidade	%	Entidade	%
Bolsonaro	12,32	PT	6,68	Brasil	15,64
Lula	9,93	Flamengo	4,03	EUA	3,15
Ciro	1,73	Palmeiras	2,47	Argentina	1,67
Dilma	0,89	STF	2,32	Rússia	1,47
Neymar	0,88	Corinthians	1,88	Europa	1,40
Moro	0,82	TSE	1,13	China	1,39
Haddad	0,62	Vasco	1,02	Ucrânia	1,24
Putin	0,59	Globo	0,99	Portugal	0,92
Jesus	0,54	Inter	0,85	BR	0,86
Rapaz	0,39	MBL	0,82	São Paulo	0,83

TABLE VI

ENTIDADES PREDOMINANTES EM COMENTÁRIOS CLASSIFICADOS COMO GATILHO, ORDENADAS POR PORCENTAGEM DE OCORRÊNCIA DENTRO DO TOTAL DE CADA CATEGORIA.

leve tendência a manter o conteúdo semântico estabelecido na discussão.

Reconhecimento de Entidades Nomeadas. De forma mais geral, nossos resultados mostram que comentários de árvores *Não Tóxicas* apresentam uma maior diversidade de presença de entidades nomeadas, combinando, por exemplo, discussões em relação ao cenário eleitoral de 2022 e entretenimento, como discussões relacionadas ao futebol. Já as árvores *Tóxicas* centralizam as discussões em uma forte polarização política, decorrente da disputa presidencial. Neste contexto, a frequência relativa de figuras centrais como Lula e Bolsonaro, aumenta expressivamente nas árvores *Tóxicas*, entre ≈ 10 –13% em comparação às menções em árvores *Não Tóxicas* (6%).

Considerando os gatilhos, a Tabela VI apresenta as dez entidades mais mencionadas nestes comentários. Embora existam menções a times de futebol, comentários gatilho são mais focados nos cenários da política e justiça (PT, STF, TSE). Adicionalmente, termos associados a conflitos ideológicos ou figuras polêmicas (Moro, MBL) também estão entre as principais entidades citadas. Curiosamente, a categoria de Localidades (LOC) nos gatilhos mantém um padrão geopolítico similar ao restante do conjunto de comentários que compõem o grupo de árvores *Tóxicas*, com destaque para a polarização nacional (Brasil) e internacional (EUA, China, Rússia), sugerindo que a discussão sobre política externa também atua como ponto de toxicidade.

Etiquetagem de classe de palavra (Pos Tagging). Árvores

Árvores Não Tóxicas		Árvores Tóxicas	
Bigrama	Freq. (%)	Bigrama	Freq. (%)
boa sorte	0,43	ser human	0,23
parte maior	0,19	classe média	0,11
ano passado	0,18	direita extr.	0,09
vez primeira	0,12	lavagem cer.	0,08
dia feliz	0,12	p*ta m*rda	0,05
ensino médio	0,12	povo bras.	0,05
coisa melhor	0,10	gente burra	0,04
salário mín.	0,10	mesma m*rda	0,04
vez última	0,10	p*u duro	0,04
prazo longo	0,09	p*u pequeno	0,04

TABLE VII

COMPARATIVO DOS 10 BIGRAMAS (SUBSTANTIVO + ADJETIVO OU VICE-VERSA) MAIS FREQUENTES DE CADA GRUPO.

Tóxicas apresentam a predominância um pouco maior de Substantivos (+0,64% mais ocorrências que nas árvores *Não Tóxicas*). Já árvores *Não Tóxicas* utilizam uma quantidade maior de Advérbios (+0,85

Ao avaliar bigramas mais frequentes e combinações de rótulos presentes a partir deles, destacamos a combinação de bigrama das etiquetas *Substantivo + Adjetivo* (e vice-versa).¹¹ Como mostrado na Tabela VII, podemos observar uma diferenciação no domínio semântico dos comentários: enquanto árvores *Não Tóxicas* são caracterizadas por construções voltadas à polidez, bem-estar e marcação temporal (como *boa sorte*, *dia feliz*, *ano passado*), árvores *Tóxicas* enfatizam a polarização ideológica e direcionam xingamentos e palavrões. Bigramas como *lavagem cerebral*, *direita extrema* e *gente burra* demonstram que, em discussões tóxicas, há mais indícios de estruturas construídas para afetar outros usuários ou demarcar polarização política.

Ao analisarmos comentários-gatilho, nossos resultados mostram uma maior predominância de conectivos de subordinação e marcadores de opinião, trigramas como *eu acho que*, *não sei se* e conjuntos de bigramas de etiquetas do tipo *Verbo + Conjunção Subordinativa*. Isso sugere que a toxicidade não é iniciada necessariamente por insultos, mas por expressões de dúvida ou discordância.

Classificação de Tipo de Discurso de Ódio. No conjunto de árvores de discussão *Tóxicas* foram identificados 12.669 comentários em que o modelo obteve pelo menos 0.5 de probabilidade de atribuição a pelo menos uma categoria de discurso de ódio. Estes comentários estão distribuídos em 7.881 árvores, representando 16,47% do total de árvores analisadas. A pequena porcentagem de árvores com comentários rotulados pelo modelo possivelmente se deve à sua dificuldade em classificar com alto grau de certeza textos informais, com

¹¹Outras combinações de bigrama foram analisadas, sendo o resultado descrito o de maior relevância para a nossa análise.

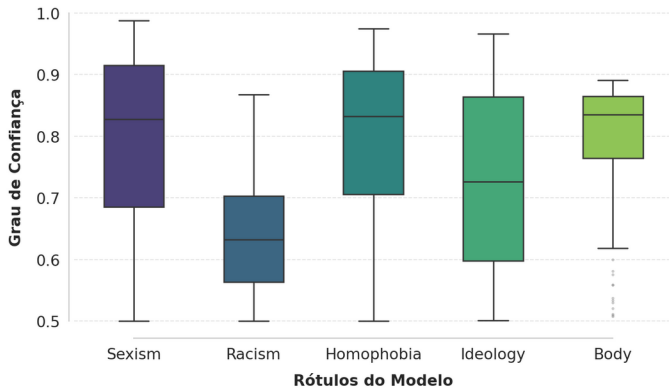


Fig. 1. Dispersão dos graus de confiança por categoria de discurso de ódio.

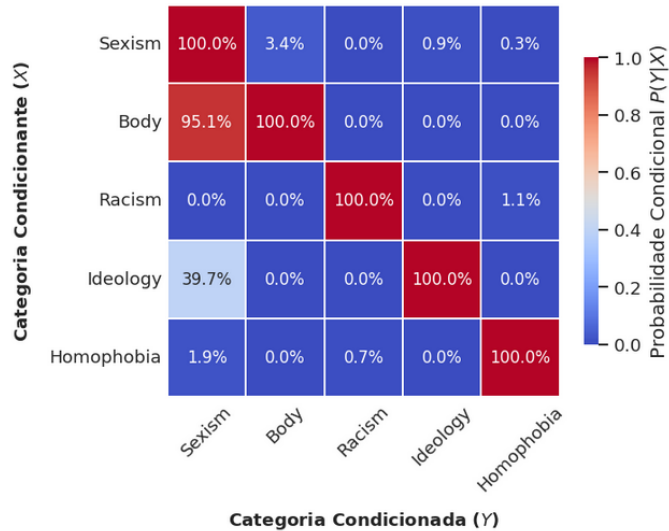


Fig. 2. Matriz de co-ocorrência condicional ($P(Y|X)$) entre as categorias de discurso de ódio.

uso intensivo de gírias e com especificidades contextuais dos eventos ocorridos no ano de 2022.

Ao analisar a distribuição da confiança do modelo para valores acima do limiar 0.5, como mostrado na Figura 1, categorias *Body*, *Homophobia* e *Sexism* apresentam as maiores medianas de confiança (≈ 0.83), indicando maior assertividade do modelo para este tipo de conteúdo. Em contrapartida, a categoria *Racism* registra a menor mediana do conjunto (≈ 0.63), sugerindo que as manifestações racistas no *corpus* tendem a ser mais sutis, estruturalmente ambíguas, ou infrequentes, dificultando a atribuição de altas probabilidades pelo classificador. Assim, o rótulo predominante mais frequente em nossos dados é *Sexism*, com 10.176 comentários, superando consideravelmente os demais rótulos: *Homophobia* com 1.351, *Racism* com 828, *Ideology* com 162 e *Body* com 152 comentários.

A seguir, analisamos os padrões de co-ocorrência entre múltiplos rótulos atribuídos pelo modelo (com limiar ≥ 0.5) a um dado comentário. A Figura 2 apresenta a matriz de

co-ocorrência condicional entre as categorias analisadas. Ela adota uma abordagem probabilística direcional ($P(Y|X)$), que calcula a probabilidade de um comentário pertencer à categoria da coluna (Y) dado que já foi classificado na categoria da linha (X). Diferentemente da correlação linear, essa análise revela assimetrias estruturais da toxicidade encontrada nos comentários. A categoria *Sexism*, por exemplo, atua como uma categoria dominante, absorvendo quase integralmente os ataques à aparência física: 95,1% dos comentários classificados como *Body* são simultaneamente sexistas. Isso sugere que, neste corpus, as críticas estéticas funcionam instrumentalmente como vetores de discurso sexista. Ressaltamos que essas tendências de interdependência se mantêm consistentes ao se observar especificamente as categorias dos comentários gatilhos.

V. LIMITAÇÕES

As limitações deste estudo, são as seguintes: (i) o desbalanceamento de classes no treinamento do modelo de detecção de toxicidade: apesar das técnicas de mitigação empregadas, a disparidade nas métricas de classe ainda é considerável e, embora uma validação humana das árvores de comentários classificadas como tóxicas fosse ideal para atestar a qualidade da detecção, essa alternativa mostrou-se inviável devido ao alto custo e tempo de anotação frente ao grande volume de dados processados; (ii) o viés político decorrente do recorte temporal da coleta de dados em 2022: por se tratar de um ano eleitoral de alta polarização, o léxico de comentários tóxicos avaliado está frequentemente atrelado ao debate político, como evidenciado na caracterização linguística deste trabalho, o que pode restringir a generalização do modelo para discursos de ódio em contextos mais cotidianos; e, por fim, (iii) o escopo restrito às dez comunidades analisadas: essa delimitação reduz a diversidade de assuntos e de estruturas de discussão avaliadas, uma vez que cada *subreddit* apresenta particularidades.

VI. CONCLUSÃO

Este trabalho caracteriza a dinâmica discursiva do Reddit brasileiro com objetivo de distinguir as estruturas de árvores de discussões *Tóxicas* e *Não Tóxicas* considerando o impacto do fenômeno de gatilhos. Assim, elaboramos um ajuste fino ao modelo BERTaBaporu, a partir de dados rotulados manualmente e técnicas de *data augmentation*, e o aplicamos à totalidade dos comentários. A análise, fundamentada no conteúdo das árvores de discussão, integrou métricas de similaridade vetorial, morfossintáticas e semânticas para caracterizar a arquitetura das interações e a incidência de discurso de ódio.

Nossos resultados mostram que árvores *Tóxicas* apresentam maior profundidade e engajamento, com coesão semântica inicial que diminui em interações mais profundas. Nestas, predominam a polarização política e a sintaxe argumentativa, em contraste com a estabilidade e neutralidade temática das árvores *Não Tóxicas*. Adicionalmente, existe uma predominância de discurso de ódio sexista em discussões *Tóxicas*. Uma vez que o tipo de árvore é definida pela presença ou não de um gatilho, nossos resultados indicam que de fato

um gatilho influencia o conteúdo, estrutura e intensidade das discussões presentes no nosso conjunto de dados.

Nosso trabalho é um primeiro passo importante para fundamentar o desenvolvimento de mecanismos preditivos para identificar gatilhos, permitindo antecipar e mitigar a toxicidade com base nos padrões identificados em discussões realizadas em comunidades brasileiras no Reddit. Como trabalhos futuros, planejamos ampliar o escopo do estudo para incluir novos *subreddits* e atualizar a base de dados com discussões mais recentes, além de analisar individualmente cada comunidade e suas potenciais particularidades de estrutura conversacional.

Considerações éticas. Para a anotação manual, divulgamos apenas o texto do comentário, o *subreddit* e a data de publicação, preservando o anonimato típico dos usuários do Reddit. Metadados de *id* serviram exclusivamente para a reconstrução das árvores de discussão.

REFERENCES

- [1] D. Vrontis, E. Siachou, G. Sakka, S. Chatterjee, R. Chaudhuri, and A. Ghosh, "Societal effects of social media in organizations: Reflective points deriving from a systematic literature review and a bibliometric meta-analysis," *European Management Journal*, vol. 40, no. 2, pp. 151–162, 2022. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0263237322000196>
- [2] M. Avallé *et al.*, "Persistent interaction patterns across social media platforms and over time," *Nature*, vol. 628, no. 8008, pp. 582–589, 2024.
- [3] M. Maleki, M. Arani, E. Mead, J. Kready, and N. Agarwal, "Applying an epidemiological model to evaluate the propagation of toxicity related to covid-19 on twitter," 01 2022.
- [4] A. Arif, M. A. Qadir, R. S. Martins, and H. M. A. Khuwaja, "The impact of cyberbullying on mental health outcomes amongst university students: A systematic review," *PLOS Mental Health*, vol. 1, no. 6, pp. 1–16, 11 2024. [Online]. Available: <https://doi.org/10.1371/journal.pmen.0000166>
- [5] N. Proferes, N. Jones, S. Gilbert, C. Fiesler, and M. Zimmer, "Studying reddit: A systematic overview of disciplines, approaches, methods, and ethics," *Social Media+ Society*, vol. 7, no. 2, 2021.
- [6] H. Almerexhi, H. Kwak, and B. J. Jansen, "Investigating toxicity changes of cross-community redditors from 2 billion posts and comments," *PeerJ Computer Science*, vol. 8, p. e1059, 2022.
- [7] Y. Yu, J. Jiang, and P. Dhillon, "Characterizing the structure of online conversations across reddit," 2024. [Online]. Available: <https://arxiv.org/abs/2209.14836>
- [8] H. Almerexhi, H. Kwak, J. Salminen, and B. J. Jansen, "Are these comments triggering? predicting triggers of toxicity in online discussions," in *Proceedings of The Web Conference 2020*, ser. WWW '20. New York, NY, USA: Association for Computing Machinery, 2020, p. 3033–3040. [Online]. Available: <https://doi.org/10.1145/3366423.3380074>
- [9] H. Almerexhi, H. Kwak, J. Salminen, and B. J. Jansen, "Provoke: Toxicity trigger detection in conversations from the top 100 subreddits," *Data and Information Management*, vol. 6, no. 4, p. 100019, 2022. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S2543925122001176>
- [10] G. Piorino, L. Lima, A. Pagano, and A. Silva, "Toxicidade e gatilhos: Um estudo de caso em comunidades do reddit no brasil," in *Proceedings of the 31st Brazilian Symposium on Multimedia and the Web*. Porto Alegre, RS, Brasil: SBC, 2025, pp. 575–579. [Online]. Available: <https://sol.sbc.org.br/index.php/webmedia/article/view/38012>
- [11] D. Hiaeshutter-Rice and I. Hawkins, "The language of extremism on social media: An examination of posts, comments, and themes on reddit," *Frontiers in Political Science*, vol. 4, p. 805008, 2022.
- [12] D. Kumar, J. Hancock, K. Thomas, and Z. Durumeric, "Understanding the behaviors of toxic accounts on reddit," in *Proceedings of the ACM Web Conference 2023*, ser. WWW '23. New York, NY, USA: Association for Computing Machinery, 2023, p. 2797–2807. [Online]. Available: <https://doi.org/10.1145/3543507.3583522>
- [13] G. Cunha and A. Silva, "Caracterizando polarização em redes sociais: Um estudo de caso das discussões no reddit sobre as eleições brasileiras de 2018 e 2022," in *Proceedings of the 30th Brazilian Symposium on Multimedia and the Web*. Porto Alegre, RS, Brasil: SBC, 2024, pp. 365–369. [Online]. Available: <https://sol.sbc.org.br/index.php/webmedia/article/view/30333>
- [14] C. Eberhart, L. Ignaczak, and M. Martins, "Text mining for cyberbullying detection: a brazilian portuguese evaluation," in *Anais do XIII Simpósio Brasileiro de Tecnologia da Informação e da Linguagem Humana*. Porto Alegre, RS, Brasil: SBC, 2021, pp. 92–100. [Online]. Available: <https://sol.sbc.org.br/index.php/stil/article/view/17788>
- [15] F. Oliveira, V. Reis, and N. Ebecken, "Tupy-e: detecting hate speech in brazilian portuguese social media with a novel dataset and comprehensive analysis of models," 2023. [Online]. Available: <https://arxiv.org/abs/2312.17704>
- [16] V. Shankaran and R. Sharma, "Analyzing toxicity in deep conversations: A reddit case study," 2024. [Online]. Available: <https://arxiv.org/abs/2404.07879>
- [17] L. H. Q. Lima, A. S. Pagano, and A. P. C. da Silva, "Toxic content detection in online social networks: a new dataset from Brazilian Reddit communities," in *Proceedings of the 16th International Conference on Computational Processing of Portuguese - Vol. 1*, P. Gamallo, D. Claro, A. Teixeira, L. Real, M. Garcia, H. G. Oliveira, and R. Amaro, Eds. Santiago de Compostela, Galicia/Spain: Association for

- Computational Linguistics, Mar. 2024, pp. 472–482. [Online]. Available: <https://aclanthology.org/2024.propor-1.48/>
- [18] P. B. Costa *et al.*, “BERTabaporu: Assessing a genre-specific language model for Portuguese NLP,” in *Proceedings of the 14th International Conference on Recent Advances in Natural Language Processing*. INCOMA Ltd., Shoumen, Bulgaria, 2023, pp. 217–223.
- [19] T. Wolf *et al.*, “Transformers: State-of-the-art natural language processing,” in *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*. Association for Computational Linguistics, 2020, pp. 38–45.
- [20] S. Bird, E. Klein, and E. Loper, *Natural Language Processing with Python: Analyzing Text with the Natural Language Toolkit*. O’Reilly Media, Inc., 2009.
- [21] J. A. Leite *et al.*, “Toxic language detection in social media for Brazilian Portuguese: New dataset and multilingual analysis,” in *Proceedings of the 1st Conference of the Asia-Pacific and the 10th International Joint Conference on Natural Language Processing*. Suzhou, China: Association for Computational Linguistics, 2020, pp. 914–924.
- [22] F. Vargas *et al.*, “HateBR: A large expert annotated corpus of Brazilian Instagram comments for offensive language and hate speech detection,” in *Proceedings of the Thirteenth Language Resources and Evaluation Conference*. Marseille, France: European Language Resources Association, 2022, pp. 7174–7183.
- [23] T. S. Almeida, H. Abonizio, R. Nogueira, and R. Pires, “Sabiá-2: A new generation of portuguese large language models,” 2024. [Online]. Available: <https://arxiv.org/abs/2403.09887>
- [24] H. Touvron *et al.*, “Llama 2: Open foundation and fine-tuned chat models,” 2023. [Online]. Available: <https://arxiv.org/abs/2307.09288>
- [25] A. Q. Jiang *et al.*, “Mistral 7b,” 2023. [Online]. Available: <https://arxiv.org/abs/2310.06825>
- [26] N. Reimers and I. Gurevych, “Sentence-bert: Sentence embeddings using siamese bert-networks,” in *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing*. Association for Computational Linguistics, 11 2019. [Online]. Available: <http://arxiv.org/abs/1908.10084>
- [27] M. A. Hearst, “Text tiling: Segmenting text into multi-paragraph subtopic passages,” *Computational Linguistics*, vol. 23, no. 1, pp. 33–64, 1997. [Online]. Available: <https://aclanthology.org/J97-1003/>
- [28] J. Nothman, N. Ringland, W. Radford, T. Murphy, and J. R. Curran, “Learning multilingual named entity recognition from wikipedia,” *Artificial Intelligence*, vol. 194, pp. 151–175, 2013.
- [29] S. Petrov, D. Das, and R. McDonald, “A universal part-of-speech tagset,” *arXiv preprint arXiv:1104.2086*, 2011.
- [30] spaCy, “Portuguese models,” <https://spacy.io/models/pt>, 2023, Acessado em: 09/11/2025.
- [31] A. Rademaker, F. Chalub, L. Real, C. Freitas, E. Bick, and V. de Paiva, “Universal dependencies for portuguese,” in *Proceedings of the Fourth International Conference on Dependency Linguistics (Depling)*, Pisa, Italy, September 2017, pp. 197–206. [Online]. Available: <http://aclweb.org/anthology/W17-6523>
- [32] C. Freitas, P. Rocha, and E. Bick, “A new world in floresta sintá(c)tica – the portuguese treebank,” *Calidoscópico*, vol. 6, no. 3, pp. 142–148, 2008. [Online]. Available: <https://revistas.unisinos.br/index.php/calidoscopio/article/view/5256/2510>
- [33] G. Sidorov, F. Velasquez, E. Stamatatos, A. Gelbukh, and L. Chanona-Hernández, “Syntactic n-grams as machine learning features for natural language processing,” *Expert Systems with Applications*, vol. 41, no. 3, pp. 853–860, 2014.
- [34] P. Fortuna, J. Rocha da Silva, J. Soler-Company, L. Wanner, and S. Nunes, “A hierarchically-labeled Portuguese hate speech dataset,” in *Proceedings of the Third Workshop on Abusive Language Online*. Florence, Italy: Association for Computational Linguistics, Aug. 2019, pp. 94–104. [Online]. Available: <https://aclanthology.org/W19-3510>
- [35] Pablo Botton da Costa, “bertabaporu-base-uncased (revision 1982d0f),” 2022. [Online]. Available: <https://huggingface.co/pablocosta/bertabaporu-base-uncased>
- [36] T. Davidson, D. Warmesley, M. Macy, and I. Weber, “Automated hate speech detection and the problem of offensive language,” in *Proceedings of the international AAAI conference on web and social media*, vol. 11, no. 1, 2017, pp. 512–515.
- [37] B. Lewandowska-Tomaszczyk, A. Baczkowska, C. Liebeskind, G. Valunaite Oleskeviciene, and S. Žitnik, “An integrated explicit and implicit offensive language taxonomy,” *Lodz Papers in Pragmatics*, vol. 19, no. 1, pp. 7–48, 2023.
- [38] J. S. Park, J. Seering, and M. S. Bernstein, “Measuring the prevalence of anti-social behavior in online communities,” 2022. [Online]. Available: <https://arxiv.org/abs/2208.13094>