

Estado da publicação: O preprint não foi publicado em outro meio.

MENSURANDO EMOÇÃO E MORALIDADE EM CAMPANHAS DIGITAIS: DESAFIOS METODOLÓGICOS DA ANÁLISE AUTOMATIZADA DE SENTIMENTOS EM ELEIÇÕES

Denise Clara Santos Santana

<https://doi.org/10.1590/SciELOPreprints.17870>

Submetido em: 2026-09-01

Postado em: 2026-09-03 (versão 1)

(AAAA-MM-DD)

MENSURANDO EMOÇÃO E MORALIDADE EM CAMPANHAS DIGITAIS: DESAFIOS METODOLÓGICOS DA ANÁLISE AUTOMATIZADA DE SENTIMENTOS EM ELEIÇÕES

SANTANA, Denise Clara Santos

ORCID: <https://orcid.org/0009-0009-6315-9725>

<denisesantanapro45@gmail.com>

Universidade Federal de Goiás. Goiânia, Goiás (GO), Brasil

RESUMO: O crescimento das redes sociais ampliou as possibilidades de estudar comportamento político a partir de grandes volumes de dados textuais, mas a operacionalização de conceitos como emoção e moralidade segue sendo um desafio metodológico. Este trabalho pergunta em que medida técnicas automatizadas de análise de texto, entre elas os dicionários léxicos, conseguem captar conteúdos moralizados e emocionalmente carregados em comunicação política digital. O material reúne 2.716 publicações originais das contas de Lula, Bolsonaro e Ciro Gomes no Twitter, nas eleições de 2022. A moralidade foi medida por dicionário de fundamentos morais, a valência emocional por classificador supervisionado, e modelos de regressão relacionam essas características ao engajamento recebido. Os indicadores reproduzem os padrões da literatura internacional, com mais difusão para conteúdo negativo e moralizado. O exame das medidas encontra dois problemas que nenhum teste convencional detecta. O resultado moral é carregado por uma única categoria do dicionário, em que 61% das ocorrências são as palavras bom e boa. O agrupamento temático separou o corpus por formato de publicação antes de separá-lo por assunto, e é esse controle que elimina a associação da emoção. O trabalho organiza esses achados em uma tipologia e conclui que a validade de uma medida depende da relação entre o instrumento e o corpus.

Palavras-chave: validade de mensuração, análise automatizada de conteúdo, fundamentos morais, comunicação política digital, eleições de 2022

MEASURING EMOTION AND MORALITY IN DIGITAL CAMPAIGNS: METHODOLOGICAL CHALLENGES OF AUTOMATED SENTIMENT ANALYSIS IN ELECTIONS

ABSTRACT: The growth of social media has expanded the possibilities of studying political behaviour from large volumes of textual data, but operationalizing concepts such as emotion and morality remains a methodological challenge. This paper asks to what extent automated text analysis techniques, lexical dictionaries among them, can capture moralized and emotionally charged content in digital political communication. The material comprises 2,716 original Twitter posts by Lula, Bolsonaro and Ciro Gomes in the 2022 Brazilian elections. Morality was measured with a moral foundations dictionary, emotional valence with a supervised classifier, and regression models relate these features to the engagement received. The indicators reproduce the patterns of the international literature, with greater diffusion for negative and moralized content. Examining the measures reveals two problems that no conventional test detects. The moral result is carried by a single dictionary category, in which 61% of its occurrences are the words bom and boa. Topic modeling split the corpus by post format before splitting it by subject, and it is this control that eliminates the association of emotion. The paper organizes these findings into a typology and concludes that the validity of a measure depends on the relation between instrument and corpus.

Keywords: measurement validity, automated content analysis, moral foundations, digital political communication, 2022 brazilian elections

INTRODUÇÃO

A comunicação política produz hoje um registro público e contínuo de mensagens e das reações que elas provocam. Nas plataformas digitais, cada publicação pode ser observada junto às interações que recebe, tornando possível investigar, em grandes volumes de dados, por que alguns conteúdos circulam mais do que outros. Esse novo ambiente ampliou as possibilidades de estudar a difusão das mensagens, mas também tornou mais urgente uma questão anterior: como mensurar, em escala, os atributos dos textos que supostamente explicam sua circulação?

Parte da literatura do campo tem atribuído papel importante à emoção e à moralidade. Brady et al. (2017) encontraram, em um corpus de tweets sobre temas políticos controversos, que a presença de linguagem moral-emocional aumentava a probabilidade de difusão de uma mensagem, padrão mantido pela replicação pré-registrada e pela meta-análise posteriores do grupo (BRADY et al., 2025). Em uma síntese mais ampla, Van Bavel et al. (2024) relacionam esses resultados à própria dinâmica da comunicação política nas

plataformas ao passo em que sua arquitetura tenderia a favorecer sistematicamente conteúdos carregados de julgamento moral e afeto negativo.

Testar afirmações desse tipo, contudo, exige transformar conceitos como emoção e moralidade em medidas observáveis. É nessa passagem que se concentra o problema examinado neste trabalho. Emoção e moralidade são conceitos construídos por tradições de pesquisa que, durante muito tempo, dependeram de leitura próxima, interpretação e consideração do contexto. Aplicá-los a centenas de milhares de textos exige instrumentos capazes de classificar grandes volumes de dados de maneira rápida e consistente. Todo instrumento de mensuração em escala, porém, impõe uma troca: ganha-se capacidade de processamento, mas pode-se perder sensibilidade ao significado do texto.

A questão, portanto, não consiste apenas em identificar palavras associadas a determinado conceito. Uma medida pode apresentar associação estatística com um desfecho e, ainda assim, representar de forma inadequada o fenômeno que pretende medir. Indicadores automatizados de emoção e moralidade não observam diretamente emoção ou moralidade; observam sinais textuais que o pesquisador toma como indicadores desses construtos. Uma associação forte entre um indicador e o número de retuítes demonstra que o indicador se relaciona ao desfecho, mas não estabelece, por si só, que ele represente adequadamente a aplicação teórica. Em pesquisas de análise automatizada de texto, essa distinção é especialmente importante porque a escala da classificação pode ocultar erros de interpretação que seriam evidentes em uma leitura manual de casos individuais.

Medidas baseadas em dicionários apresentam correlação modesta com avaliações humanas em tarefas de análise psicológica de texto (RATHJE et al., 2024), e o uso direto de rótulos automatizados como variáveis em modelos de regressão pode produzir viés e intervalos de confiança inadequados mesmo quando a acurácia da classificação é elevada (EGAMI et al., 2023). Esses estudos não apontam para a inviabilidade da mensuração automatizada. O que mostram é que o erro do instrumento precisa ser conhecido e avaliado no contexto em que a medida é utilizada.

É desse problema que deriva a pergunta deste trabalho: em que medida indicadores automatizados de emoção e moralidade reproduzem os padrões de associação com a difusão documentados na literatura e, ao mesmo tempo, preservam validade de mensuração em um corpus eleitoral brasileiro?

Para examinar essa questão, o estudo reúne 2.716 publicações originais das contas oficiais de Luiz Inácio Lula da Silva, Jair Bolsonaro e Ciro Gomes no Twitter, publicadas

entre 16 de agosto e 30 de outubro de 2022 e recuperadas a partir dos identificadores do Interfaces Twitter Elections Dataset (IASULAITIS et al., 2025). A moralidade foi operacionalizada por meio da contagem de termos de um dicionário léxico (CARVALHO et al., 2020), enquanto a valência emocional foi estimada por um classificador supervisionado treinado em português brasileiro. As duas medidas foram incorporadas a regressões lineares tendo como variável dependente o logaritmo do número de retuítes.

O desenho é exploratório. Não se pretende estimar o erro dos instrumentos, o que exigiria anotação humana, nem estabelecer qual estratégia de operacionalização é superior, o que exigiria aplicar as duas ao mesmo construto. O que se faz é aplicar os procedimentos usuais do campo a um caso concreto e examinar, resultado por resultado, o que os indicadores permitem e o que não permitem afirmar sobre os conceitos que dizem representar.

Os resultados produzem uma resposta em duas etapas. De um lado, os indicadores reproduzem associações compatíveis com aquelas documentadas pela literatura internacional. De outro, a análise revela problemas de mensuração que impedem interpretar essas associações como evidência suficiente de validade dos construtos. No caso da moralidade, o resultado é fortemente dependente de uma única categoria do dicionário, cujas ocorrências correspondem, em 61% dos casos, às palavras bom e boa. A associação entre intensidade emocional e difusão, por sua vez, aparece ou desaparece conforme a especificação do controle temático. Além disso, a medida de moralidade perde capacidade de discriminação em textos muito curtos.

O trabalho tem quatro seções, além desta introdução e das considerações finais. A primeira reconstrói o que a literatura afirma sobre emoção, moralidade e difusão, e de que maneira ela transforma esses conceitos em números. A segunda descreve o corpus, os instrumentos e o modelo. A terceira apresenta os resultados e submete cada um deles ao exame da medida que o produziu. A quarta organiza os problemas encontrados em uma tipologia.

1 EMOÇÃO, MORALIDADE E DIFUSÃO: O QUE A LITERATURA MEDE

A literatura sobre difusão de mensagens em redes sociais atribui papel crescente à emoção e à moralidade como propriedades capazes de ampliar sua circulação. Brady et al. (2017), em análise de 563.312 tweets sobre temas políticos polêmicos, encontraram que cada palavra moral-emocional acrescentada a uma mensagem aumentava em cerca de 20% a probabilidade de compartilhamento. A replicação pré-registrada e a meta-análise

posteriores do grupo confirmaram esse padrão (BRADY et al., 2025). Outros trabalhos identificam associações semelhantes para diferentes dimensões da linguagem emocional, moral e identitária, sugerindo que conteúdos que mobilizam afetos, julgamentos morais e distinções entre grupos tendem a alcançar maior circulação (SCHÖNE; PARKINSON; GOLDENBERG, 2021; WU et al., 2025).

Apesar da convergência dos achados, esses estudos não transformam os conceitos em números da mesma maneira. Uma primeira estratégia é a utilização de dicionários, nos quais determinados termos são previamente associados a uma categoria e sua ocorrência é convertida em uma medida quantitativa. Brady et al. (2017), por exemplo, identificam linguagem moral-emocional por meio da contagem de palavras de um dicionário. Rathje, Van Bavel e van der Linden (2021) utilizam termos referentes ao adversário político como indicador de linguagem de exogrupo. No estudo de Kramer, Guillory e Hancock (2014), o LIWC desempenha função semelhante ao classificar palavras segundo categorias emocionais positivas e negativas. Nessas medidas, a classificação depende da correspondência entre o vocabulário encontrado no texto e uma lista previamente definida.

Uma segunda estratégia utiliza classificadores treinados a partir de exemplos anotados. Em vez de especificar antecipadamente quais palavras representam determinado conceito, o pesquisador fornece ao modelo textos previamente classificados e utiliza os padrões aprendidos para produzir classificações em novos casos. Antypas, Preece e Camacho-Collados (2023), por exemplo, empregam modelos de linguagem para classificar tweets de parlamentares. A diferença é substantiva já que enquanto o dicionário estabelece a medida a partir de um vocabulário definido antes da análise, o classificador procura inferir a classificação a partir de regularidades observadas nos dados de treinamento que são supervisionados por seres humanos.

Uma terceira estratégia aparece quando o fenômeno de interesse é observado fora do próprio texto. Soroka, Fournier e Nir (2019), em estudo psicofisiológico com 1.156 participantes em dezessete países, mediram a reação corporal a conteúdos noticiosos positivos e negativos. Nesse caso, a negatividade não é inferida pela presença de determinadas palavras; ela é relacionada à resposta produzida pelo indivíduo diante do conteúdo. Quando a mesma questão é investigada em plataformas digitais, contudo, a mensuração frequentemente retorna ao texto, como nos estudos que contam palavras negativas ou classificam mensagens por modelos automatizados. O contraste evidencia que medir uma propriedade do texto e medir a reação provocada pelo texto são operações distintas, ainda que ambas possam ser utilizadas para testar hipóteses sobre difusão.

Essas diferenças importam porque cada estratégia estabelece uma relação diferente entre conceito, indicador e evidência. Um dicionário pressupõe que determinados termos sejam bons representantes do conceito; um classificador pressupõe que os exemplos utilizados em seu treinamento contendo informação suficiente para distinguir as categorias; uma medida comportamental ou fisiológica pressupõe que a resposta observada constitua um indicador válido do fenômeno teórico. O fato de métodos diferentes produzirem associações semelhantes com a difusão, portanto, não significa necessariamente que estejam medindo exatamente o mesmo fenômeno.

No caso da moralidade, essa questão permanece. Brady, Crockett e Van Bavel (2020) sistematizam, no modelo MAD (Motivation, Attention, and Design), mecanismos pelos quais conteúdos moralizados podem favorecer atenção e difusão. A evidência empírica mobilizada por essa literatura, contudo, frequentemente operacionaliza a moralidade por meio de indicadores lexicais. A presença e a frequência de termos associados a fundamentos morais passam a representar a intensidade da moralização de uma mensagem. Esse procedimento permite transformar grandes corpus em variáveis comparáveis, mas introduz a questão sobre em que medida a ocorrência desses termos representa efetivamente a presença e a intensidade da moralidade no texto.

A literatura metodológica recente mostra que esse problema não é apenas teórico. Rathje et al. (2024), analisando quinze conjuntos de dados e 47.925 mensagens anotadas manualmente em doze idiomas, encontraram correlações entre 0,20 e 0,30 entre medidas baseadas em dicionários e avaliações humanas. Egami et al. (2023) demonstram, adicionalmente, que utilizar classificações automatizadas diretamente como variáveis em modelos de regressão pode produzir viés e intervalos de confiança inadequados, mesmo quando a acurácia do classificador é elevada.

A meta-análise de Brady et al. (2025) reforça a importância dessa questão ao registrar heterogeneidade entre estudos, com diferenças nos tamanhos de efeito segundo plataforma, idioma e período analisado. Parte dessa variação pode refletir diferenças substantivas entre os contextos estudados. Outra parte pode decorrer da própria mensuração, uma vez que instrumentos lexicais construídos para idiomas e contextos diferentes não necessariamente capturam os mesmos conceitos com a mesma precisão.

No contexto brasileiro, essa questão permanece pouco explorada. Carvalho et al. (2020) construíram e validaram o MFD-BR, versão em português brasileiro do dicionário de fundamentos morais, que vem sendo aplicado ao discurso político. Mundim e

Mendonça (2026) identificaram com esse instrumento a arquitetura moral do discurso de Bolsonaro, com predominância de autoridade e lealdade, e Dantas (2023) o utilizou para acompanhar deslocamentos morais no discurso eleitoral de candidatos de direita. Em língua portuguesa há ainda um dicionário para o português europeu aplicado a debates parlamentares (ZÁQUETE; ORGHIAN; PINHEIRO, 2023) e um estudo comparativo entre Brasil e Portugal sobre argumentos morais no debate climático do Twitter (COSTA; CAPOANO; BALBÉ, 2022). Cristiani, Lieira e Camargo (2020) demonstraram a viabilidade da análise de sentimento em corpus eleitoral brasileiro, enquanto D'Ignazi, Kalimeri e Beiró (2025) combinaram análise de sentimento e marcadores morais em notícias.

2 DADOS E PROCEDIMENTOS

Esta seção descreve o caso sobre o qual os instrumentos foram aplicados. Convém deixar explícito, de saída, o que o desenho não permite. As duas estratégias descritas na seção anterior não foram cruzadas com os dois conceitos, porque o dicionário mede a moralidade e o classificador mede a emoção. A diferença entre os instrumentos está, portanto, confundida com a diferença entre os construtos, e nada do que se segue compara o desempenho de um com o do outro. O que a aplicação conjunta permite é observar onde cada instrumento falha na tarefa que lhe coube.

2.1 O corpus

O material provém do Interfaces Twitter Elections Dataset, publicado por Iasulaitis et al. (2025), que reúne identificadores de publicações da eleição presidencial brasileira de 2022 e cobre o período de 20 de julho de 2022 a 23 de fevereiro de 2023. Em conformidade com os termos de serviço da plataforma, o conjunto é distribuído sem os textos, e cabe a cada pesquisador recuperá-los a partir dos identificadores numéricos.

A janela analisada vai de 16 de agosto a 30 de outubro de 2022, do início da propaganda eleitoral ao segundo turno. A unidade de análise é o tweet original, definido pela ausência de referência a outra publicação, o que deixa de fora as continuções de sequência, os tweets que citam outros e as respostas do próprio candidato, porque uma continuação herda a audiência do primeiro tweet e sua contagem de retuítes não mede o apelo do próprio texto. Dos 3.546 identificadores da janela, 3.515 tiveram o texto recuperado, 2.717 eram publicações originais e 2.716 tinham texto não vazio, que é o corpus analítico. Ele reúne 1.606 publicações de Lula, 291 de Bolsonaro e 819 de Ciro

Gomes. A assimetria de volume é um dado da campanha, porque Bolsonaro publicou no Twitter, na janela analisada, cerca de um quinto do que publicou Lula, padrão compatível com a ênfase de sua campanha em outras plataformas (CHICARINO; CONCEIÇÃO; ALVES, 2024).

2.2 A medida moral

A moralidade é medida com o MFD-BR (CARVALHO et al., 2020), que reúne cerca de setecentas entradas distribuídas pelos cinco fundamentos da Teoria dos Fundamentos Morais e por uma categoria adicional de moralidade geral. Três propriedades explicam a adoção desse tipo de instrumento. Ele é transparente, porque a lista é pública e qualquer resultado se rastreia até os termos que o produziram. Ele é barato, porque a contagem roda em segundos sobre milhões de textos. Ele é estável, porque a mesma lista aplicada ao mesmo corpus devolve sempre o mesmo número. Nenhuma das três propriedades pertence ao classificador.

A categoria de moralidade geral merece descrição à parte. Ela reúne 66 das setecentas entradas e não corresponde a nenhum dos cinco fundamentos. O que seus termos têm em comum é sinalizar avaliação moral sem indicar qual fundamento está em jogo, e a lista vai do julgamento explícito, com moral, imoral, íntegro, perverso e transgredir, à avaliação genérica, com bom, boa, correto, ruim e legal. É a única das seis categorias sem contraparte na Teoria dos Fundamentos Morais, e a seção 3.2 mostra o que isso produz quando ela entra no modelo.

As limitações derivam do mesmo ponto, o de a contagem tratar o texto como conjunto de palavras sem ordem. A negação inverte o sentido de uma sentença sem alterar a contagem. A ironia inverte o sentido sem alterar nenhuma palavra. A citação do adversário coloca no texto do candidato um vocabulário que não é dele. Há ainda um problema de cobertura, porque uma lista construída para um domínio carrega as escolhas desse domínio, e o MFD-BR foi construído para classificação de notícias falsas (CARVALHO et al., 2020).

A aplicação segue o critério de presença. Para cada fundamento atribui-se valor 1 quando a publicação contém ao menos um termo da categoria e 0 quando não contém, sem peso distinto para virtude e vício.

Os cinco fundamentos foram agregados em três indicadores. Uso de fundamentos individualizantes vale 1 quando a publicação contém termo de cuidado ou de justiça, e uso

de fundamentos vinculantes vale 1 quando contém termo de lealdade, autoridade ou santidade, separação que Graham, Haidt e Nosek (2009) estabeleceram ao comparar campos ideológicos. Uso de moralidade geral vale 1 quando a publicação contém termo da categoria de moralidade do dicionário, que não corresponde a nenhum fundamento específico. No corpus, 15,2% das publicações acionam fundamentos individualizantes, 25,3% acionam vinculantes e 12,8% apresentam vocabulário de moralidade geral. Considerando ao menos um dos três, 41,2% das publicações apresentam alguma linguagem moral.

2.3 A medida emocional

A valência emocional vem de um classificador de sentimento da biblioteca pysentimiento (PÉREZ; GIUDICI; LUQUE, 2021), construído sobre o BERTweet.BR, modelo de linguagem pré-treinado em tweets em português brasileiro (CARNEIRO et al., 2025). Para cada texto ele devolve três probabilidades, correspondentes às categorias positivo, negativo e neutro, cuja soma é igual a um. Diferente do dicionário, ele processa a sentença inteira e leva em conta a posição de cada palavra, o que lhe permite tratar negação e construções em que o sentido depende da ordem. O Quadro 1 reúne os termos da comparação entre os dois instrumentos.

Quadro 1 — Os dois instrumentos utilizados

Critério	Dicionário léxico	Classificador supervisionado
Unidade lida	a palavra isolada	a sentença inteira, com a ordem das palavras
Origem da validade	juízo de quem montou a lista	textos anotados por humanos no treinamento
Rastreabilidade	total, o termo que gerou o resultado é identificável	nenhuma, o resultado não se decompõe em termos
Custo de aplicação	segundos para milhões de textos	horas para dezenas de milhares em máquina comum
Estabilidade	idêntica entre execuções	sensível à versão da biblioteca
Trata negação e ironia	não	em parte
Conceito medido neste trabalho	moralidade	valência emocional

Fonte: elaboração própria. Os dois instrumentos não foram aplicados ao mesmo construto, conforme a ressalva da abertura desta seção.

Das três probabilidades devolvidas pelo classificador, são construídas duas variáveis. A intensidade emocional mede o afastamento da neutralidade e a assimetria de valência mede qual polo predomina:

$$\text{Intensidade} = 1 - P(\text{neutro})$$

$$\text{Assimetria} = P(\text{negativo}) - P(\text{positivo})$$

A intensidade varia de 0, na neutralidade completa, a 1, na carga emocional máxima, e não distingue a direção do afeto. No corpus tem média de 0,446 e desvio-padrão de

0,314. A assimetria varia de menos um, quando predomina inteiramente o conteúdo positivo, a mais um, quando predomina inteiramente o negativo. No corpus tem média de $-0,074$ e desvio-padrão de $0,487$.

2.4 O controle temático

Os grupos temáticos não foram definidos previamente a partir de uma lista de categorias. Eles resultaram de modelagem de tópicos com o BERTopic (GROOTENDORST, 2022), procedimento que representa cada publicação por um vetor numérico, reduz a dimensão desse vetor e agrupa as publicações próximas. A aplicação ao corpus produziu os dez grupos da Tabela 1, usados como variável de controle. A natureza desses grupos é examinada na seção 3.3.

Tabela 1 — Grupos temáticos produzidos pela modelagem de tópicos

Grupo	Termos que o distinguem	Publicações	Mediana de retuítes
0	brasil, povo, brasildaesperança, país, vamos	960	1.124
1	bolsonaro, povo, presidente, pra, porque	637	1.360
2	governo, renda, salário, vai, mínimo	232	1.213
3	educação, casa, universidades, vamos, ensino	165	1.011
4	ciro tv, vivo, cirotv br, br, acompanhe	180	384
5	broadcasts, lula, imprensa, conversa, caminhada	139	1.003
6	hoje, ciro, dia, amanhã, povo	170	1.826
7	ciro, giro ciro, giro, perca, sp	122	3.628
8	bolsonaro, saúde, vacina, pandemia, popular	63	1.619
9	indígenas, povos, precisamos, preservação, vamos	48	1.085

Fonte: elaboração própria. $N = 2.716$. Os termos são os cinco de maior peso em cada grupo. A numeração é a que o procedimento atribuiu, e a natureza dos grupos é examinada na seção 3.3.

2.5 O modelo

A distribuição do número de retuítes é assimétrica. Nenhuma publicação tem valor zero, a menor contagem é 139 e poucos casos concentram valores altos. Na escala original a assimetria é 3,49, e depois da transformação logarítmica cai para 0,50. A variável dependente é, portanto, o logaritmo do número de retuítes, estimado por regressão linear por mínimos quadrados com erros-padrão robustos. Como a variável dependente está em logaritmo, cada coeficiente exprime variação proporcional, e a conversão exata é a exponencial do coeficiente menos um.

O modelo explica o logaritmo dos retuítes pelas propriedades do conteúdo e pelos controles:

$$\begin{aligned} \log(\text{retuítes}) = & \beta_0 + \beta_1 \text{ Intensidade} + \beta_2 \text{ Assimetria} \\ & + \beta_3 \text{ Individualizantes} + \beta_4 \text{ Vinculantes} + \beta_5 \text{ Moralidade geral} \end{aligned}$$

$$+ (\text{Individualizantes} + \text{Vinculantes}) \times \text{Candidato} \\ + \text{Candidato} + \text{Dia de evento} + \text{Tema} + \text{Semana} + \text{erro}$$

O coeficiente β_1 mede a associação entre intensidade emocional e difusão, e β_2 a da assimetria de valência. Os coeficientes β_3 , β_4 e β_5 medem, cada um, a diferença associada ao uso de um dos três tipos de linguagem moral, já descontado o que a intensidade e a valência explicam. Os termos de interação verificam se essa associação varia conforme o emissor. Tema e Semana são blocos de indicadores binários, um para cada um dos dez grupos temáticos e um para cada um dos onze blocos semanais. O tempo entra por semana, e não por data. O modelo tem 32 parâmetros, 2.716 observações e coeficiente de determinação ajustado de 0,669.

3 O FUNCIONAMENTO APARENTE DOS INDICADORES

Aplicados a este corpus, os indicadores geram associações estatisticamente significativas e alinhadas com a literatura internacional, que é o que se espera de uma medida em funcionamento. Essa validade, porém, é aparente. A Tabela 2 traz os coeficientes de interesse.

Tabela 2 — Associação entre conteúdo e logaritmo dos retuítes

Variável	Coefficiente	Efeito	p
Intensidade emocional (β_1)	0,0702	+7,3%	0,148
Assimetria de valência (β_2)	0,1640	+17,8%	< 0,001 ***
Usa fundamentos individualizantes (β_3)	-0,0282	-2,8%	0,397
Usa fundamentos vinculantes (β_4)	-0,0100	-1,0%	0,730
Usa moralidade geral (β_5)	0,1537	+16,6%	< 0,001 ***
Dia de evento de campanha	0,1040	+11,0%	0,009 **

Fonte: elaboração própria. N = 2.716. Regressão linear por mínimos quadrados com erros-padrão robustos, 32 parâmetros, coeficiente de determinação ajustado de 0,669. Maior fator de inflação de variância de 1,08 e teste de especificação de Ramsey com p de 0,405. Controles de candidato, grupo temático, semana de campanha e dia de evento omitidos da tabela. Asteriscos indicam significância estatística, com *** para p abaixo de 0,001, ** para p abaixo de 0,01 e * para p abaixo de 0,05.

3.1 A negatividade

A assimetria de valência é a associação mais bem estabelecida do modelo. Quanto mais a mensagem se desloca para o polo negativo, mais ela circula. A escala vai de menos um a mais um, e cada unidade percorrida corresponde a 17,8% a mais de retuítes, o que dá cerca de 8% para uma diferença de um desvio-padrão. O padrão aparece em desenhos bastante diferentes entre si, da reação psicofisiológica medida por Soroka, Fournier e Nir (2019) ao consumo de notícias examinado por Robertson et al. (2023) e ao corpus multilíngue de Twitter analisado por Antypas, Preece e Camacho-Collados (2023).

3.2 A moralidade

A Tabela 2 traz três indicadores de linguagem moral, e apenas um deles se sustenta. Publicações com vocabulário de moralidade geral recebem 16,6% a mais de retuítes, com p abaixo de 0,001. Os dois eixos substantivos da Teoria dos Fundamentos Morais não aparecem. O individualizante fica em $-2,8\%$ com p de 0,397, e o vinculante em $-1,0\%$ com p de 0,730, ambos negativos e indistinguíveis de zero.

Entretanto, o peso está todo na categoria que menos se parece com a teoria. Um levantamento das ocorrências dessa categoria no corpus mostra que 61% delas são as palavras bom e boa. Dizer que publicações com linguagem moral recebem 16,6% a mais de retuítes é, em boa parte, dizer que publicações com as palavras bom ou boa recebem mais retuítes. Como descrição do corpus, o achado continua verdadeiro. Como evidência ele não se sustenta, porque avaliação positiva difusa não é julgamento moral no sentido da teoria.

A causa é a cobertura do dicionário, que falha nos dois sentidos. Ele inclui termos de uso majoritariamente não moral, como bom e legal, e não inclui termos que carregam julgamento moral inequívoco no português da campanha, como corrupto e vergonha. As duas falhas acrescentam ruído à medida e retiram dela sinal. Nenhum teste de significância separa uma coisa da outra, porque o que ele mede é se a associação existe, e não se a variável mede o conceito. O teste de interação entre candidato e fundamento moral, derivado da literatura de congruência ideológica (FEINBERG; WILLER, 2015), produz um padrão contrário ao esperado, que também não permite distinguir efeito substantivo de artefato de mensuração.

3.3 A intensidade emocional

A intensidade emocional não se associa à difusão. O coeficiente é positivo, mas o p fica em 0,148, o que não permite distingui-lo de zero. Isso contraria a previsão mais direta da literatura, segundo a qual carga emocional de qualquer direção circularia mais. A causa não está no classificador que mede a emoção. Está na variável de controle temático, e é por isso que este é o achado metodológico mais importante do trabalho. Incluir ou não esse controle muda a resposta.

O controle veio de uma modelagem de tópicos, que agrupou as publicações por semelhança de vocabulário sem receber categorias prontas. Dos dez grupos da Tabela 1, quatro correspondem a assuntos no sentido usual, que são os grupos 2, 3, 8 e 9, de economia, educação, saúde e povos indígenas. Outros quatro correspondem a formatos de

publicação, que são os grupos 4, 5, 6 e 7, de transmissões ao vivo, registros de agenda e saudações, e reúnem 22,5% do corpus. Os grupos 0 e 1 concentram 58,8% do corpus e trazem linguagem geral de campanha. A separação por formato não é interpretação, e foi conferida no texto. No grupo de transmissões de atos e caminhadas, dominado por Lula, 78% das publicações trazem marcadores explícitos como ao vivo, assista ou o endereço da transmissão, e no grupo de transmissões e canais próprios, dominado por Ciro Gomes, são 62%. No resto do corpus esses marcadores aparecem em 3%.

Esses grupos de formato são a origem do problema, porque estão ligados ao engajamento e à emoção ao mesmo tempo. A publicação mediana dos dois grupos de transmissão recebe 741 retuítes, contra 1.250 das demais, e a intensidade emocional média deles é de 0,172, contra 0,482 das demais. São publicações protocolares, que circulam menos e quase não carregam carga emocional. Sem o controle, elas produzem sozinhas uma correlação entre emoção e retuítes que não vem de emoção nenhuma. Com o controle, a associação some e o p sobe para 0,148. A diferença de intensidade se mantém dentro de cada faixa de comprimento de texto, então não é efeito de os anúncios serem curtos, e trocar os indicadores de semana por indicadores de data move o p só para 0,094.

A lição ultrapassa este corpus. Quem aplicasse o mesmo desenho sem controle temático concluiria que a intensidade emocional está associada à difusão. Quem usasse uma lista de temas definida de antemão classificaria os anúncios de transmissão pelo assunto de que eles tratam e chegaria à mesma conclusão. Só o agrupamento automático, que não recebeu categoria alguma, expôs a divisão por formato. A vantagem, neste caso, foi o procedimento não partir de categorias definidas de antemão.

4 UMA TIPOLOGIA DOS PROBLEMAS DE MENSURAÇÃO

Os problemas encontrados na seção anterior não são do mesmo tipo, e a diferença entre eles determina o que cada um exigiria para ser resolvido. Dois estão em pontos identificáveis do desenho. O primeiro está dentro da própria medida, quando a variável reúne termos que não representam o conceito. Foi o caso da categoria de moralidade geral do dicionário, descrita na seção 2.2, cujas ocorrências no corpus são em 61% dos casos as palavras bom e boa. O segundo está na variável de controle, quando o procedimento que deveria isolar o efeito de interesse separa o corpus por outro critério. Foi o caso do agrupamento temático, que dividiu as publicações por formato antes de dividi-las por assunto.

A esses dois somam-se duas condições que não se prendem a um resultado específico e afetam todos eles. O Quadro 2 reúne os quatro casos e registra, na coluna da direita, o que fica em aberto em cada um. A distinção entre os dois grupos está no alcance. Nos dois primeiros, o que se discute é a interpretação de um resultado que o modelo já produziu. Nos dois últimos, o que se discute é se o resultado se sustenta fora das condições exatas em que foi obtido.

Quadro 2 — Tipologia dos problemas de mensuração encontrados

Tipo	Onde aparece neste corpus	O que fica em questão
Validade de conteúdo (3.2)	a categoria de moralidade geral, em que 61% das ocorrências são bom e boa	o que a variável de fato mede
Validade do controle (3.3)	o agrupamento temático separa por formato de publicação, e não por assunto	qual conclusão o modelo devolve
Confiabilidade dos dados (4.1)	a resposta de erro da plataforma na coleta das respostas	de que corpus se está falando
Validação ausente (4.2)	nenhum instrumento foi checado por leitura humana	o tamanho do erro, que segue desconhecido

Fonte: elaboração própria. Os dois primeiros são achados deste corpus e dizem respeito à validade das medidas. Os dois últimos são condições que atravessam todos os resultados apresentados.

4.1 Confiabilidade dos dados

A confiabilidade dos dados não é garantida pela plataforma. No corpus complementar de respostas, recuperado quase quatro anos depois do pleito, 86% das requisições voltaram com código de erro, o que levaria a estimar que essa proporção das respostas havia sido apagada. Em teste sobre 600 identificadores nunca requisitados antes, 515 deram erro na primeira passada e 431 deles, ou 83,7%, devolveram a página completa quando a requisição foi repetida em seguida. O código media o intervalo entre requisições, e não a exclusão de conteúdo, e a taxa real fica em torno de 14%. Antes desse diagnóstico, cinco taxas diferentes foram produzidas e descartadas, cada uma sustentando uma explicação plausível. O caso vale para qualquer corpus recuperado de plataforma, inclusive o principal deste trabalho. Freelon (2018) descreve o acesso restrito como a condição de trabalho do período posterior às interfaces abertas, e aqui se acrescenta que o obstáculo não é só de acesso, porque a própria resposta pode ser lida errado.

4.2 Limitações

A limitação principal do trabalho é a ausência de checagem contra leitura humana. Nem o classificador nem o dicionário foram avaliados sobre este corpus, e a validade de ambos é herdada dos trabalhos que os construíram, nenhum deles sobre discurso eleitoral brasileiro. Uma correlação entre 0,20 e 0,30 com a avaliação humana, como a que Rathje et al. (2024) medem para dicionários, tratada como se fosse medida exata em uma regressão,

pode atenuar os coeficientes e produzir intervalos de confiança estreitos demais, que é o problema de inferência formalizado por Egami et al. (2023).

O procedimento que enfrentaria isso é conhecido. Ele começa por um livro de códigos com definição operacional de cada categoria, na forma que Halterman e Keith (2026) recomendam, e segue para uma amostra estratificada por candidato, por período e por extensão de texto, codificada por dois avaliadores independentes com concordância estimada. A comparação entre essa codificação e a do instrumento automatizado fornece a estimativa de erro usada na correção. A amostra precisaria conter negação, ironia e citação do adversário, raras no corpus e frequentes no erro, que uma amostra aleatória simples sub-representa. Esse procedimento não foi executado aqui.

CONSIDERAÇÕES FINAIS

Este trabalho perguntou em que medida indicadores automatizados de emoção e moralidade reproduzem os padrões de associação com a difusão documentados na literatura e, ao mesmo tempo, preservam validade de mensuração. A resposta obtida sobre 2.716 publicações da eleição brasileira de 2022 tem duas partes que precisam ser mantidas juntas.

A primeira parte é que as técnicas captam padrões empíricos compatíveis com resultados documentados em outros contextos. Elas recuperam a associação entre negatividade e difusão, com 17,8% a mais de retuítes por unidade da escala de assimetria, e sustentam a associação entre linguagem moral e difusão depois de controlada a emoção. O custo computacional de tudo isso coube em uma máquina pessoal sem placa gráfica dedicada, o que torna esse tipo de análise acessível a pesquisas.

A segunda parte é que reproduzir esses padrões não estabelece a validade das medidas. O resultado moral, com p abaixo de 0,001, vem de uma categoria cujas ocorrências são em 61% dos casos as palavras bom e boa. A associação entre emoção e difusão desaparece ou permanece conforme a inclusão de um controle temático cuja natureza só se percebe lendo os grupos que o procedimento produziu. Nenhuma das duas coisas aparece em um valor de p , em um coeficiente de determinação ou em um teste de pressuposto.

A lição principal deste trabalho é metodológica. O que os dados dizem sobre a campanha de 2022 depende inteiramente do que os instrumentos estavam medindo, e é justamente isso que permanece por estabelecer. A validade de uma medida de conteúdo

não é propriedade do instrumento, e sim da relação entre o instrumento, o conceito e o corpus em que ele é aplicado.

A tipologia da seção 4 serve a esse fim, ao separar os problemas que estão na medida e no controle das condições de confiabilidade e de validação que atravessam qualquer resultado. A separação importa porque as soluções são diferentes. Validade de conteúdo se enfrenta decompondo e inspecionando a medida, e daí a recomendação de publicar a decomposição dos indicadores agregados, sem a qual o problema da categoria de moralidade geral seria invisível. Validade do controle se enfrenta lendo o que o procedimento automático produziu, em lugar de tratá-lo como caixa fechada. Nenhuma delas se enfrenta com mais verificações de robustez sobre a mesma medida.

Sendo exploratório, o trabalho não estabelece o tamanho do erro de nenhum dos instrumentos, e a agenda que ele abre é justamente essa. A anotação humana de uma amostra estratificada, com estimativa explícita do erro de classificação, é o passo que converteria as perguntas levantadas aqui em discussões mais elaboradas.

A contribuição principal deste trabalho é mostrar que a validade de uma medida automatizada de conteúdo não pode ser pressuposta a partir da significância estatística nem da reprodução de padrões conhecidos. Ela exige o exame da composição da medida, da natureza dos controles e das condições em que o instrumento é aplicado, e esse exame produz achados que nenhum teste de robustez convencional é capaz de detectar.

MATERIAIS SUPLEMENTARES

A análise usou dois conjuntos de dados. O primeiro reúne as publicações das contas oficiais dos três candidatos, recuperadas a partir dos identificadores do Interfaces Twitter Elections Dataset (IASULAITIS et al., 2025), disponível em <https://doi.org/10.1371/journal.pone.0316626> sob licença CC BY-NC-SA 4.0. O segundo reúne 48.124 respostas públicas a essas publicações, coletadas para esta pesquisa em julho de 2026, sem registro de nome de usuário, de identificador de perfil ou de endereço. Sobre esse segundo corpus foi estimada uma análise complementar, cujo resultado substantivo não integra o argumento deste paper.

Um pacote de replicação reúne o código de limpeza, classificação, agrupamento e estimação, em sete etapas que reconstroem os resultados a partir das bases, com uma etapa final que recalcula os valores declarados no texto e acusa divergência. O código é distribuído sob licença MIT e os dados derivados sob CC BY-NC-SA 4.0, herdada da

fonte. Por exigência dos termos da plataforma, as respostas circulam apenas como identificadores, sem o texto.

DECLARAÇÃO DE USO DE INTELIGÊNCIA ARTIFICIAL

Em atendimento à Portaria CNPq nº 2.664/2026, declara-se o uso de ferramentas de inteligência artificial generativa neste trabalho. Claude Opus 5 e Claude Sonnet 5 foram empregados em programação, depuração de scripts, verificação de consistência, revisão de redação e tradução. Consensus, SciSpace e Research Rabbit foram empregados na busca bibliográfica, e todas as referências assim localizadas foram consultadas em suas fontes originais. A concepção da pesquisa, as decisões metodológicas, a análise e a interpretação dos resultados constituem trabalho autoral. Os modelos de linguagem empregados como instrumentos de medida, descritos na seção 2, não integram esta declaração.

REFERÊNCIAS

ANTYPAS, Dimosthenis; PREECE, Alun; CAMACHO-COLLADOS, Jose. Negativity spreads faster: a large-scale multilingual Twitter analysis on the role of sentiment in political communication. **Online Social Networks and Media**, v. 33, art. 100242, 2023.

BRADY, William J.; CROCKETT, M. J.; VAN BAVEL, Jay J. The MAD model of moral contagion: the role of motivation, attention, and design in the spread of moralized content online. **Perspectives on Psychological Science**, v. 15, n. 4, p. 978-1010, 2020.

BRADY, William J.; RATHJE, Steve; GLOBIG, Laura K.; VAN BAVEL, Jay J. Estimating the effect size of moral contagion in online networks: a pre-registered replication and meta-analysis. **PNAS Nexus**, v. 4, n. 11, art. pgaf327, 2025.

BRADY, William J.; WILLS, Julian A.; JOST, John T.; TUCKER, Joshua A.; VAN BAVEL, Jay J. Emotion shapes the diffusion of moralized content in social networks. **Proceedings of the National Academy of Sciences**, v. 114, n. 28, p. 7313-7318, 2017.

CARNEIRO, Fernando; VIANNA, Daniela; CARVALHO, Jonnathan; PLASTINO, Alexandre; PAES, Aline. BERTweet.BR: a pre-trained language model for tweets in Portuguese. **Neural Computing and Applications**, v. 37, n. 6, p. 4363-4385, 2025.

CARVALHO, Flavio; OKUNO, Helder Yukio; BARONI, Lais; GUEDES, Gustavo. A Brazilian Portuguese Moral Foundations Dictionary for fake news classification. In: INTERNATIONAL CONFERENCE OF THE CHILEAN COMPUTER SCIENCE SOCIETY (SCCC), 39., 2020, Coquimbo. **Proceedings** [...]. [S. l.]: IEEE, 2020.

CHICARINO, Tathiana Senne; CONCEIÇÃO, Desirée Luíse Lopes; ALVES, Márcia. Engajamento e populismo: a presença digital de Lula e Bolsonaro nas eleições 2022 no Twitter. **Política & Trabalho: Revista de Ciências Sociais**, n. 60, p. 105-124, 2024.

COSTA, Pedro Rodrigues; CAPOANO, Edson; BALBÉ, Alice. Alterações climáticas e argumentos morais no Twitter: um estudo comparativo entre Brasil e Portugal. **Interações: Sociedade e as Novas Modernidades**, n. 43, 2022.

CRISTIANI, André; LIEIRA, Douglas; CAMARGO, Heloisa A. A sentiment analysis of Brazilian elections tweets. In: SYMPOSIUM ON KNOWLEDGE DISCOVERY,

MINING AND LEARNING (KDMiLe), 2020, Porto Alegre. **Anais [...]**. Porto Alegre: SBC, 2020. p. 153-160.

DANTAS, Leonardo Feitosa. **Electoral discourse of right-wing candidates in redistributionist districts: a study of changing moral foundations**. 2023. Dissertação (Mestrado) - Escola Brasileira de Administração Pública e de Empresas, Fundação Getúlio Vargas, Rio de Janeiro, 2023.

D'IGNAZI, Jacopo; KALIMERI, Kyriaki; BEIRÓ, Mariano G. **Beyond sentiment: examining the role of moral foundations in user engagement with news on Twitter**. [S. l.]: arXiv, 2025. arXiv:2502.12009.

EGAMI, Naoki; HINCK, Musashi; STEWART, Brandon M.; WEI, Hanying. Using imperfect surrogates for downstream inference: design-based supervised learning for social science applications of large language models. **Advances in Neural Information Processing Systems**, v. 36, p. 68589-68601, 2023.

FEINBERG, Matthew; WILLER, Robb. From gulf to bridge: when do moral arguments facilitate political influence? **Personality and Social Psychology Bulletin**, v. 41, n. 12, p. 1665-1681, 2015.

FRELON, Deen. Computational research in the post-API age. **Political Communication**, v. 35, n. 4, p. 665-668, 2018.

GRAHAM, Jesse; HAIDT, Jonathan; NOSEK, Brian A. Liberals and conservatives rely on different sets of moral foundations. **Journal of Personality and Social Psychology**, v. 96, n. 5, p. 1029-1046, 2009.

GROOTENDORST, Maarten. **BERTopic: neural topic modeling with a class- based TF-IDF procedure**. [S. l.]: arXiv, 2022. arXiv:2203.05794.

HALTERMAN, Andrew; KEITH, Katherine A. Codebook LLMs: evaluating LLMs as measurement tools for political science concepts. **Political Analysis**, v. 34, n. 2, p. 188-204, 2026.

IASULATTIS, Sylvia; VALEJO, Alan Demétrius Baria; GRECO, Bruno Cardoso; PERILLO, Vinicius Gonçalves; MESSIAS, Guilherme Henrique; VICARI, Isabella. The Interfaces Twitter Elections Dataset: construction process and characteristics of big social data during the 2022 presidential elections in Brazil. **PLOS ONE**, v. 20, n. 2, art. e0316626, 2025.

KRAMER, Adam D. I.; GUILLORY, Jamie E.; HANCOCK, Jeffrey T. Experimental evidence of massive-scale emotional contagion through social networks. **Proceedings of the National Academy of Sciences**, v. 111, n. 24, p. 8788-8790, 2014.

MUNDIM, Pedro Santos; MENDONÇA, Marcella Soares. O protetor e o justiceiro: a arquitetura moral do populismo de direita nos discursos de Bolsonaro (2018) e Milei (2023). **Revista Debates**, Porto Alegre, v. 20, n. 1, p. 64-82, jan./abr. 2026.

PÉREZ, Juan Manuel; GIUDICI, Juan Carlos; LUQUE, Franco. **pysentimiento: a Python toolkit for opinion mining and social NLP tasks**. [S. l.]: arXiv, 2021. arXiv:2106.09462.

RATHJE, Steve; MIREA, Dan-Mircea; SUCHOLUTSKY, Ilia; MARJIEH, Raja; ROBERTSON, Claire E.; VAN BAVEL, Jay J. GPT is an effective tool for multilingual psychological text analysis. **Proceedings of the National Academy of Sciences**, v. 121, n. 34, art. e2308950121, 2024.

RATHJE, Steve; VAN BAVEL, Jay J.; VAN DER LINDEN, Sander. Out-group animosity drives engagement on social media. **Proceedings of the National Academy of Sciences**, v. 118, n. 26, art. e2024292118, 2021.

ROBERTSON, Claire E.; PRÖLLOCHS, Nicolas; SCHWARZENEGGER, Kaoru; PÄRNAMETS, Philip; VAN BAVEL, Jay J.; FEUERRIEGEL, Stefan. Negativity drives online news consumption. **Nature Human Behaviour**, v. 7, p. 812-822, 2023.

SCHÖNE, Jonas Paul; PARKINSON, Brian; GOLDENBERG, Amit. Negativity spreads more than positivity on Twitter after both positive and negative political situations. **Affective Science**, v. 2, p. 379-390, 2021.

SOROKA, Stuart; FOURNIER, Patrick; NIR, Lilach. Cross-national evidence of a negativity bias in psychophysiological reactions to news. **Proceedings of the National Academy of Sciences**, v. 116, n. 38, p. 18888-18892, 2019.

VAN BAVEL, Jay J.; ROBERTSON, Claire E.; DEL ROSARIO, Kareena; RATHJE, Steve. Social media and morality. **Annual Review of Psychology**, v. 75, p. 311-340, 2024.

WU, Chunqiong; JIANG, Shan; SUN, Jianhong; LIU, Yingqi. Research on the influence mechanism of emotional communication on Twitter (X) and the effect of spreading public anger. **Acta Psychologica**, v. 260, art. 105560, 2025.

ZÁQUETE, Mafalda; ORGHIAN, Diana; PINHEIRO, Flávio L. A Moral Foundations Dictionary for the European Portuguese language: the case of Portuguese parliamentary debates. In: INTERNATIONAL CONFERENCE ON COMPUTATIONAL SCIENCE (ICCS), 2023. **Proceedings** [...]. Cham: Springer, 2023.

DECLARAÇÃO DE DISPONIBILIDADE DE DADOS DA PESQUISA

Todo o conjunto de dados de apoio aos resultados deste estudo foi publicado no artigo e na seção Materiais suplementares.

FINANCIAMENTO

O presente trabalho foi realizado com apoio da Coordenação de Aperfeiçoamento de Pessoal de Nível Superior, Brasil (CAPES), Código de Financiamento 001.

CONTRIBUIÇÃO DA AUTORA

Denise Clara Santos Santana: conceituação, curadoria de dados, análise formal, metodologia, programação, validação, visualização e redação.

DECLARAÇÃO DE CONFLITO DE INTERESSE

A autora declara não haver conflito de interesses a mencionar.

MINIBIOGRAFIA DA AUTORA

Denise Clara Santos Santana é mestranda em Ciência Política na Universidade Federal de Goiás, com bolsa da CAPES, e pesquisa comunicação política digital, emoções e moralidade em campanhas eleitorais.

Este preprint foi submetido sob as seguintes condições:

- Os autores declaram que os necessários Termos de Consentimento Livre e Esclarecido de participantes ou pacientes na pesquisa foram obtidos e estão descritos no manuscrito, quando aplicável.
- Os autores declaram que a elaboração do manuscrito seguiu as normas éticas de comunicação científica.
- Os autores declaram que estão cientes que são os únicos responsáveis pelo conteúdo do preprint e que o depósito no SciELO Preprints não significa nenhum compromisso de parte do SciELO, exceto sua preservação e disseminação.
- Os autores declaram que os dados, aplicativos e outros conteúdos subjacentes ao manuscrito estão referenciados.
- O manuscrito depositado está no formato PDF.
- Os autores declaram que a pesquisa que deu origem ao manuscrito seguiu as boas práticas éticas e que as necessárias aprovações de comitês de ética de pesquisa, quando aplicável, estão descritas no manuscrito.
- Os autores declaram que uma vez que um manuscrito é postado no servidor SciELO Preprints, o mesmo só poderá ser retirado mediante pedido à Secretaria Editorial do SciELO Preprints, que afixará um aviso de retratação no seu lugar.
- Os autores concordam que o manuscrito aprovado será disponibilizado sob licença [Creative Commons CC-BY](#).
- O autor submissor declara que as contribuições de todos os autores e declaração de conflito de interesses estão incluídas de maneira explícita e em seções específicas do manuscrito.
- Os autores declaram que o manuscrito não foi depositado e/ou disponibilizado previamente em outro servidor de preprints ou publicado em um periódico.
- Caso o manuscrito esteja em processo de avaliação ou sendo preparado para publicação mas ainda não publicado por um periódico, os autores declaram que receberam autorização do periódico para realizar este depósito.
- O autor submissor declara que todos os autores do manuscrito concordam com a submissão ao SciELO Preprints.