

PROCESSAMENTO EM LINGUAGEM NATURAL PARA ANÁLISE DE SENTIMENTOS DE ADOLESCENTES DE ESCOLA DE FUTEBOL

BRITO, Diogo de Freitas^{1,3}; BRASIL, Roxana Macedo²; BARRETO, Ana Cristina Lopes y Glória³; JUNIOR, Homero da Silva Nahum^{3,4}

Resumo

O estudo objetivou minerar a opinião de adolescentes sobre a prática em escola de futebol. Os bancos de dados de 2023 e 2024 possuíam 107 registros cada, permitindo a obtenção das 30 palavras mais recorrentes. O modelo desenvolvido em Python utilizou o método de aprendizagem de máquina, conquistando Precisão = 97,52%, Recall = 92,17%, e capacidade de discriminar sentimentos positivos e negativos de AUC-ROC = 98,37%. A ausência de diferença entre os dados de 2023 e 2024 viabilizou a submissão do último banco de dados ao modelo. Essa conquistou Precisão = 100,00%, Recall = 98,82%, F1-Score = 99,41%, e Acurácia = 99,06%. Então, possível foi concluir que o modelo atingiu satisfatoriamente o objetivo.

Palavras-chave: Inteligência artificial. Negócios. Administração. Ciência de dados. Cliente externo.

Abstract

The aim of the study was to gauge the opinions of teenagers about football practice at school. The 2023 and 2024 databases had 107 records each, allowing the 30 most recurrent words to be obtained. The model developed in Python used the machine learning method, achieving Precision = 97.52%, Recall = 92.17%, and the ability to discriminate between positive and negative sentiments of AUC-ROC = 98.37%. The lack of difference between the 2023 and 2024 data made it possible to submit the latter database to the model. It achieved Precision = 100.00%, Recall = 98.82%, F1-Score = 99.41%, and Accuracy = 99.06%. It was therefore possible to conclude that the model satisfactorily achieved its objective.

Keywords: Artificial intelligence. Business. Management. Data science. External customer.

Introdução

As aplicações de Processamento de Linguagem Natural (PLN) objetivariam a identificação de significados e representações em textos desenvolvidos em algum idioma humano. Para tanto, utilizaria estruturas gramaticais (regras e propriedades); classes de palavras (por exemplo, substantivo e adjetivo); ambiguidades; anáforas, figura de

¹ Docente do Curso de Gestão Desportiva e do Lazer do Centro Universitário Celso Lisboa;

² Docente Ph.D. em Educação Física;

³ Docentes do Curso de Educação Física do Centro Universitário Celso Lisboa;

⁴ Docente da Escola de Saúde da Universidade Cândido Mendes.

linguagem caracterizada pela repetição de palavra ou expressão no início de frases; léxico, conjunto de palavras de determinado idioma e respectivos significados; tesouro, vocabulário específico de uma área do conhecimento; e ontologia, significado de ações e entidades (Indurkha e Damerou, 2010). Portanto, a sequência de caracteres e estrutura hierárquica do idioma seriam considerados (Liu e Zhang, 2012), atenuando a complexidade potencializada pela ambiguidade (Ribeiro, Sá e Nogueira, 2024; Batista e Araújo, 2023) e pelas incorreções (Nascimento e Henz, 2021; Almeida, 2019) comumente presentes na verbalização e escrita humanas, mas elevando a exigência sobre a modelagem do banco de dados (Araújo, 2024).

Talvez, isso justifique, parcialmente, as aplicações na análise do Decreto Regulatório da Inovação (Decreto n. 9.283/2018) como esclarecedor da Lei de Inovação (Lei n. 10.973/2004), inclusive naqueles pontos considerados falhos (Quintella e Hanna, 2024); no desenvolvimento de agente conversacional com fins educacionais (Soares, Tarouco e Silva, 2021; Campos, 2002) e robótica educacional (Dos Santos, 2018); para controle defluxo de processos em negócios (Cabral, 2021); para interação homem-máquina (Lapa, 2021) e pesquisa por informações de saúde (Fernandes, 2021) em redes sociais; à avaliação de produtos pelos clientes externos (Jabeen, 2024); e ao ensino de desenvolvimento de programas de computador (Campos, 2019).

Isso somente seria possível pelo fato de que o PLN buscaria correlacionar as distinções linguísticas estratificada (sintaxe, semântica e pragmática) e em granularidade (oração versus discurso), isso seria necessário à modelagem, mas não suficiente, requisitando etapas para sustentar as distinções, assim, a partir do texto seriam realizadas *tokenização* e as análises, nessa ordem, léxica, sintática, semântica e pragmática, permitindo como desfecho a extração do significado pretendido pelo emissor (Dale, 2010).

A *tokenização* seria a segmentação de palavras, ou seja, a identificação de cada palavra, as que são denominadas de *tokens*. Nos idiomas não segmentados, as palavras seriam escritas sem separação, por exemplo no mandarim, então a *tokenização* requisita informações léxicas e morfológicas. Nos idiomas com delimitação de espaço, como o português, as palavras seriam separadas por espaços em branco, mas os *tokens* não necessariamente seriam formados pelos caracteres delimitados, pois dependeria do padrão escolhido à *tokenização* e da existência de ambiguidade no sistema de escrita, como o uso de ponto, vírgulas, aspas, apóstrofes e hifens (Palmer, 2010). Por exemplo, o ponto poderia

ser utilizado para indicar abreviação (sr.), classe numérica (1.000) ou finalização de determinada ideia.

A análise léxica consistiria em determinar a maneira como as palavras seriam trabalhadas, podendo ser como a sequência de caracteres ou forma canônica (dicionarizada). Por exemplo, *jogar*, no primeiro caso poderia ser somente a reunião ordenada das letras, ou, no segundo caso, o *lemma*, objeto abstrato, do conjunto de variantes morfológicas {jogando, jogado, jogo, jogador, jogador}. Nesse caso, a construção de dicionários de *lemmas* se faria necessário para relacioná-los às informações invariantes (semânticas e sintáticas). A etapa em questão permitiria o *parsing side*, mapeamento da palavra até o *lemma*; e no sentido contrário, a geração morfológica (*morphological generation*). A importância desses procedimentos residiria na recuperação de informação, a identificação de palavra morfológicamente complexa, exigiria o pré-processamento (*stemming*) para identificar o respectivo radical (*stem*). Exemplificando, no verbo *jogar*, o *lemma* seria *jogar*, mas o *stem* seria *jog*, pois a partir desse possível seria criar outras palavras (Hippisley, 2010).

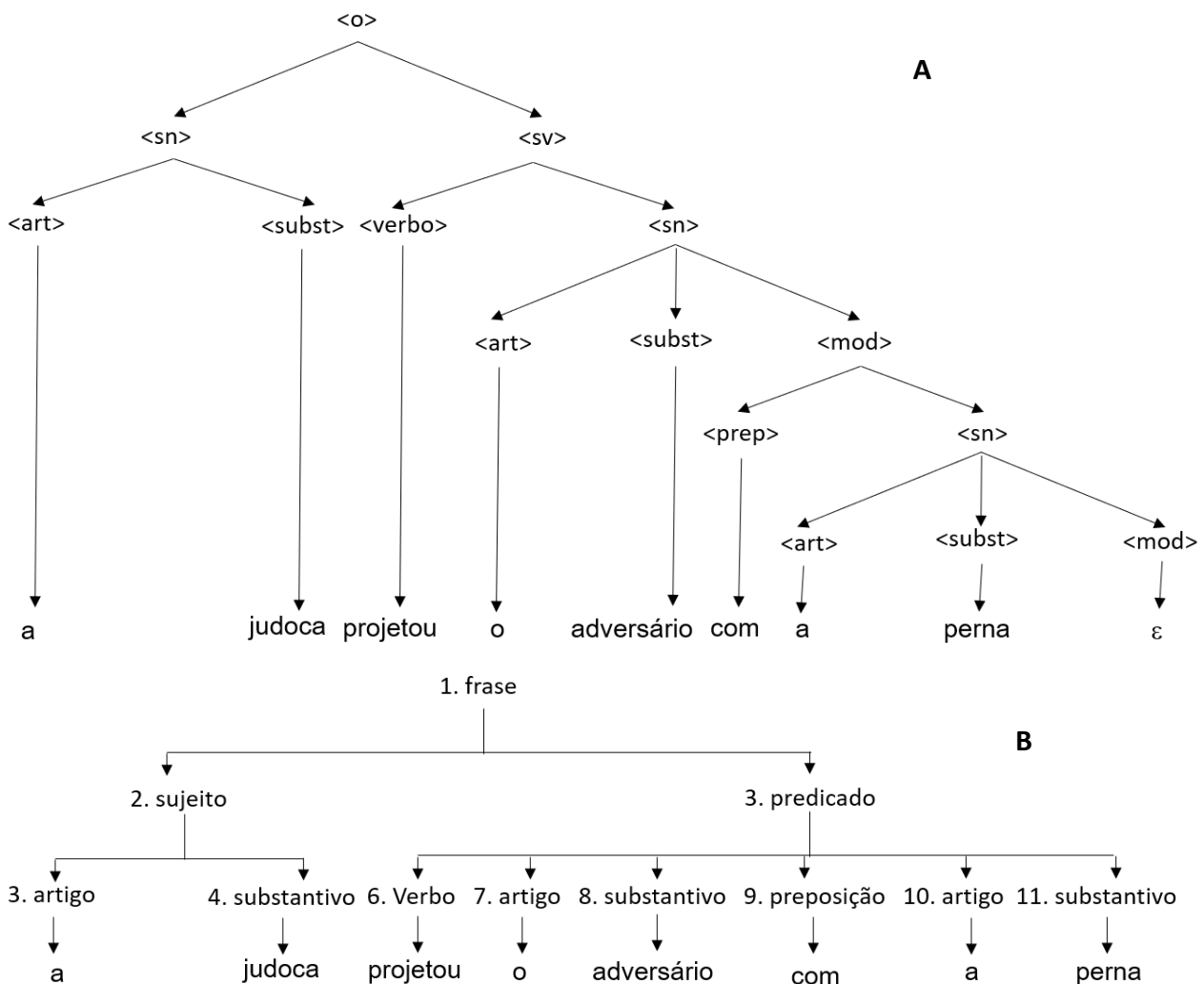
Ponto pacífico em PLN estaria na frase como a expressão de proposição, ideia ou pensamento, portanto a unidade básica de análise (sintática), requisitando a determinação da respectiva estrutura sintática ou gramatical, o que seria realizado pelo desenvolvimento de gramáticas e árvores sintáticas (*syntax tree* ou *parse tree*). Exemplificando, a construção de gramática para a frase *a judoca projetou o adversário com a perna*. O modelo Formal (Figura 1A) teria à esquerda os elementos não terminais, enquanto que à direita residiriam os símbolos terminais, ou seja, lexicais. Como a frase estaria livre de contexto, então a construção Alternativa seria satisfatória (Figura 1B). Diante do exposto, a gramática poderia ser utilizada para reconhecer e construir frases pertencentes ao idioma, processos denominados reconhecimento e geração, respectivamente. A árvore sintática representaria a plenitude das etapas de derivação de determinada frase, tal qual o desenvolvimento de árvores de decisão (Ribeiro *et al.*, 2024), em consequência, a leitura seria de cima (raiz) para baixo (ramos e folhas). Contudo, a diferença estaria na aplicação de alguma regra da gramática, compulsoriamente em cada nó. A representação poderia ser Formal (Figura 2A) ou Alternativa (Figura 2B), dependendo do contexto ou complexidade inerentes (Moura, 1996).

Figura 1: Construção Formal (A) e Alternativa (B) de uma Gramática.

A <o> → <sn><sv> <sn> → <art><subst><mod> <subst><mod> <pron> <sv> → <verbo> <verbo><sn> <verbo><sp> <verbo><sn><sp> <mod> → <adj> <sp> <sp> → <prep><sn>	B frase ⇒ sujeito predicado sujeito ⇒ artigo substantivo predicado ⇒ verbo artigo substantivo preposição artigo ⇒ a o substantivo ⇒ judoca adversário perna verbo ⇒ projetou preposição ⇒ com
--	---

Fonte: Os Autores (2025).

Figura 2: Construção Formal (A) e Alternativa (B) de uma Árvore Sintática.



Fonte: Os Autores (2025).

A compreensão da mensagem humana se classificaria como processo complexo, detendo dependência lexical, sintática, do contexto e raciocínio recorrente, essa seria a essência da análise semântica (Santos *et al.*, 2017). O complicador faria morada na origem da mensagem, pois o uso das palavras (frequência, padrão, colocação, dedução e implicação) teria característica personalíssima, logo o significado das expressões pertenceria ao domínio da metalinguagem (Anchiêta e Prado, 2022), portanto rica fonte de ambiguidade léxica (mesma palavra com significados distintos), de escopo (frase com mais de um sentido) e de referência (uso indistinto de algum termo). Nesse caso, comumente ocorreria quando o termo serviria como pronome relativo, utilizado no início de determinada oração, referenciando algum antecedente, ou conjunção integrante, introduzindo orações substantivas (Moraes *et al.*, 2025; Bunnell, 2024; Laporte, 2001).

Finalmente, a análise pragmática objetivaria reconhecer o significado de determinada frase a partir do contexto (Mioni, Barbosa e Oliveira, 2024; Bulegon e Moro, 2010; Minker, 1998). Comumente seriam utilizadas gramáticas baseadas em casos, cuja utilização se daria pela comparação de determinada frase aos padrões de construção, na existência de convergência, poderia ser considerada pertinente ao contexto (Carbonell e Hayes, 1987), o ponto sensível seria o conjunto de referências pronominais, cuja resolução disporia de diversos algoritmos, todavia dependentes da estrutura sintática (Gil, 2024; Haghighi e Klein, 2009; Vicedo González, Ferrández Rodríguez e Peral Cortés, 2000; Rocha, 2000; Rosa, 1997; Lappin e Leass, 1994).

Em suma, o PLN permitiria compreender informações textuais, o que poderia ser potencializado pela aprendizagem de máquina (*machine learning* - ML), a qual possibilitaria que sistemas “aprendessem”, fornecendo respostas melhores sem programação explícita (Arão, 2024; Ludemir, 2021). Então, na existência de léxicos associados a regras e banco de dados específicos possível seria treinar determinado algoritmo de ML para identificar sentimentos. Esses entendidos como interpretações de reações fisiológicas automáticas e circunstanciais, as emoções (Roazzi *et al.*, 2011; Fiorin, 2007).

A coadunação desses conceitos originou a análise de sentimentos ou mineração de opinião, que consistiria na identificação do teor positivo, neutro ou negativo contido no texto (Silva, 2022), aspecto particularmente relevante em negócios por possibilitar às Organizações adequações em produtos e serviços, visando melhorar ou reforçar a percepção sobre aqueles ou a marca (Alonso e Letícia, 2022; Santos, 2017), o que poderia se manifestar em (Alvarez, 2024; Freire *et al.*, 2023; Gosling *et al.*, 2019): 1) *insights*

objetivos, evitando tendenciosidades humanas; 2) desenvolvimento de produtos ou serviços (pesquisa de mercado); 3) extração de opinião em grandes volumes de dados não estruturados (*e-mails*, relacionamento com cliente externos e comentários, por exemplo); 4) personalização de resposta no atendimento ao cliente externo; e 5) acompanhamento de campanhas comerciais.

Apesar disso, não raramente, os estudos se desenvolveriam no domínio de redes sociais (Maia e Salton, 2022; Afonso e Duque, 2019; Franco e Adaniya, 2013; Lima, 2012), mesmo que considerando aspectos de saúde (Malakin, 2024; Rodas *et al.*, 2022), sociais (Antunes, Issa e Hoed, 2023), mercado de ações (Souza, Souza e Meinerz, 2021) ou metodológicos de ciências de dados (Silva e Serrano, 2023; Ribeiro e Silva, 2018).

Independentemente da aplicação, a análise de sentimentos teria as seguintes etapas algorítmicas (Kang, Yoo e Han 2012): 1) pré-processamento, identificação das palavras-chave do texto, quantificação da expressão de sentimento por pontuação entre zero (decepção) e 10 (satisfação); 2) tokenização; 3) estabelecimento do dicionário de lemas; e 4) remoção de palavras não significativas à mensagem. Com isso, a condução da pesquisa poderia utilizar método como (Liu, 2015; Rosa, Rodríguez e Bressan, 2015):

- Baseado em regras: identificaria, classificaria e pontuaria palavras-chave específicas com base nos léxicos, que receberiam a pontuação. A identificação da frase como positiva, neutra ou negativa ocorreria pela pontuação final, a qual, após ser comparada aos limites de sentimentos, geraria a carga emocional geral. Tratar-se-ia de método de fácil aplicação, porém exigiria constante retroalimentação dos léxicos e tenderia à imprecisão em razão de distinções culturais;
- ML: requisitaria o treinamento do modelo de análise de sentimentos com dados conhecidos para aplicá-lo a dados desconhecidos, cuja precisão do resultado dependeria do tamanho e da variabilidade do banco de treino, o qual deveria ser específico à área de estudo, o que poderia configurar desvantagem à ampla aplicação;
- Híbrido: resultaria da reunião das anteriores, comumente demandaria elevados investimentos de tempo, capital humano, estrutura e conhecimento.

A fragilidade da análise de sentimentos estaria em identificar e, conseqüentemente, interpretar o complexo refinamento da comunicação humana, porque exigiria a compreensão do cenário circunscritor da mensagem (Liu, 2012; Nasukawa e Yi, 2003; Dave, Lawrence e Pennock, 2003), por exemplo: 1) sarcasmo tenderia a ser interpretado

como positivo; 2) negação, quando palavras negativas a transmitir alguma inversão de significado, especialmente quando presente sequencialmente em mais de uma oração; e 3) multipolaridade, quando orações sequenciais apresentariam sentimentos distintos, mesmo com menção a aspectos diferentes, por exemplo, *gosto do livro, mas a capa é feia*. Com base no exposto, o estudo corrente objetivou minerar a opinião de adolescentes sobre a prática em escola esportiva.

Metodologia

O modelo foi desenvolvido a partir do banco de dados de uma escola de futebol localizada em Recife (PE), no qual constava a avaliação não estruturada de 107 adolescentes, dos quais 18 eram do sexo feminino, com idades entre 14 e 18 anos, regularmente matriculados e frequentadores da instituição há, pelo menos, sete meses, no ano 2023. Os participantes pertenciam a classe socioeconômica C, ou seja, com renda familiar entre três e cinco salários mínimos, conforme definido pelo Instituto Brasileiro de Geografia e Estatística (Fernandes *et al.*, 2024). O banco de dados apresentava a identificação do matriculado; frase avaliadora; pontuação, variando de 1 (insatisfeito) a 5 (satisfeito); e sentimento. Esse foi considerado positivo quando ponto = 4 ou ponto = 5, e negativo em qualquer outro caso. Assim, o total de registros positivos foi 83.

A modelagem foi desenvolvida em Python 3.10 com o emprego dos pacotes pandas 2.1.4, scikit-learn 1.5.2, nltk 3.8.1 e unidecode 1.3.8. Inicialmente, estimado foi o vetor de frequência de palavras, utilizando o método *bag of words*, o qual possibilitou identificar todos os vocábulos existentes no banco de dados e as respectivas frequências absolutas. O método à análise de sentimento foi o ML, tendo sido o banco de dados de 2023 dividido, aleatoriamente, em 70,00%, 20,00% e 10,00%, respectivamente para treino, teste e validação do modelo, garantida a representatividade pela distribuição dos sentimentos. O subconjunto de treino foi submetido à validação cruzada *k-fold*, especialmente para reduzir o risco de sobreajuste (Cunha, 2019).

A classificação (positivo e negativo) foi implementada por regressão logística, considerando a frequência relativa geral (I) e o peso de cada palavra em razão da relevância (II). Para isso, a ponderação considerou a raridade do vocábulo no banco de dados, empregando o método Frequência do Termo – Frequência Inversa nos Documentos (*Term Frequency – Inverse Document Frequency – TF-IDF*), a partir dele a relevância seria

inversamente proporcional à ocorrência (III), revelando maior expressividade do sentimento (Gelbukh *et al.*, 2024; Alrefae *et al.*, 2024; Suhasini e Vimala, 2021).

$$TF = \frac{\text{Ocorrências}}{\text{Total de palavras}} \quad (I)$$

$$IDF = \log \left(\frac{\text{Total de documentos}}{\text{Ocorrências} + 1} \right) \quad (II)$$

$$TF - IDF = TF.IDF \quad (III)$$

O contexto das declarações exigiu o emprego de N-grams, o qual possibilitaria considerar palavras consecutivas, esses conjuntos, não raramente, expressariam significados distintos daqueles dos vocábulos isolados (Alqurafi e Alsanoosy, 2024; Tashtoush *et al.*, 2024; Alvarez, Gelbukh e Sidorov, 2024), por exemplo, as palavras *jogador* e *ruim* representariam semânticas distintas e, talvez, não relacionadas, quando proferidas separadamente. Porém, a expressão *jogador ruim* deteria em si significado de qualificação do praticante, logo representando sentimento diferente daqueles quando da separação. No contexto desse estudo, empregou-se o 2-grams (bigrama), consideração de palavras aos pares.

O modelo foi aplicado ao banco de dados pareado de 2024, exigindo a investigação de discrepâncias entre os anos pelas métricas AUC-ROC (Polo e Miot, 2020), precisão, *recall* e *F1-score* (Strauss, Júnior e Ferreira, 2022). Finalmente, a distribuição dos dados de 2024 foi avaliada para garantir a ausência de mudanças (Souza, 2013).

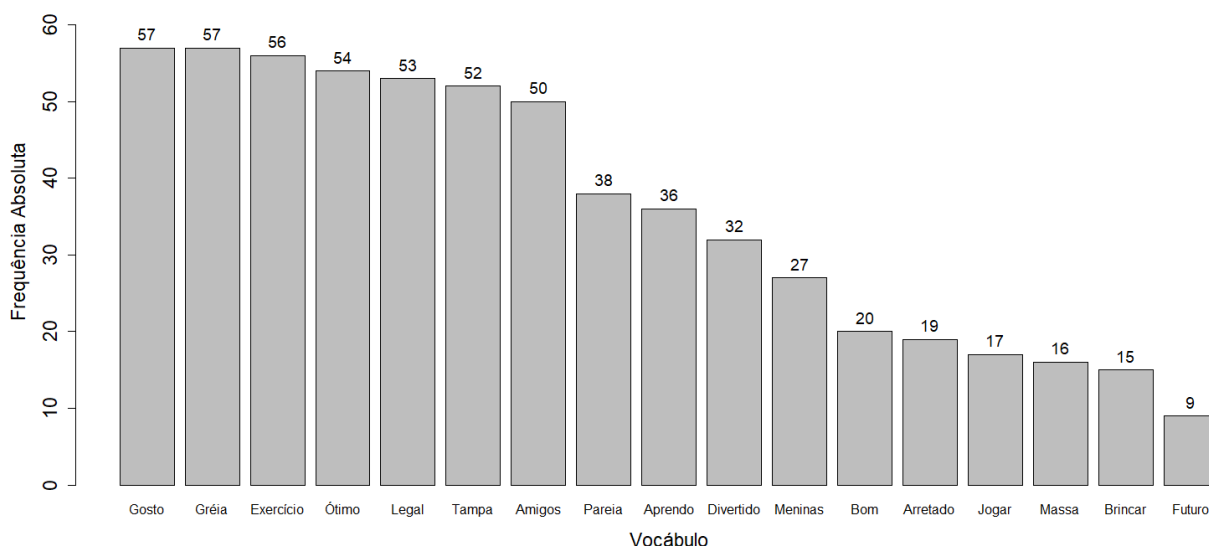
Resultados e Discussão

Inicialmente, a matriz de frequência das palavras tinha dimensão de 107 linhas e 432 colunas, convergindo ao esperado, pois contados foram artigos, preposições e pronomes, por exemplo. Todavia, aquele resultado não era parcimonioso computacionalmente e inviabilizaria a aplicação do modelo, então o *bag of words* foi refeito, porém limitando a análise de frequência às 30 palavras mais recorrentes, desconsiderando as classes mencionadas anteriormente e a pontuação por ventura existente. Essa configuração se apoiou no vocabulário médio do adolescente brasileiro, entre 10.000 e 15.000 palavras

(Alcantara *et al.*, 2021), tendo por balizamento o fenômeno de desenvolvimento lexical (Pedrosa, Dourado e Lemos, 2015), a influência das consequências tecnológicas na cognição (Silva *et al.*, 2024; Silva, 2016) e competências socioemocionais (Oliveira *et al.*, 2024).

O processo descrito possibilitou sumarizar os sentimentos pela *tokenização*, o qual requisitou processamentos sequenciais para 1) remover palavras de parada (*stop words*), 2) retirar toda pontuação, 3) reduzir as palavras ao radical (*stemming*), e 4) uniformizar a redação pela imposição de iniciais maiúsculas (Condori, 2015; Pasqualotti e Vieira, 2008). Os resultados para os vocábulos considerados Positivos (Figura 3), demonstrando empate técnico entre os sete primeiros vocábulos, e a percepção de Futuro, talvez associada à profissionalização, com apenas nove ocorrências. Valeria destacar que a expressão de sentimentos estaria associada aos valores sociais e culturais (Nepomuceno *et al.*, 2017; Budó *et al.*, 2007), característica revelada nas palavras Arretado (algo bom), Gréia (brincadeira), Massa (coisa boa), Pareia (amizade) e Tampa (alguém bom em algo).

Figura 3: Frequência dos Sentimentos Positivos após Tokenização.

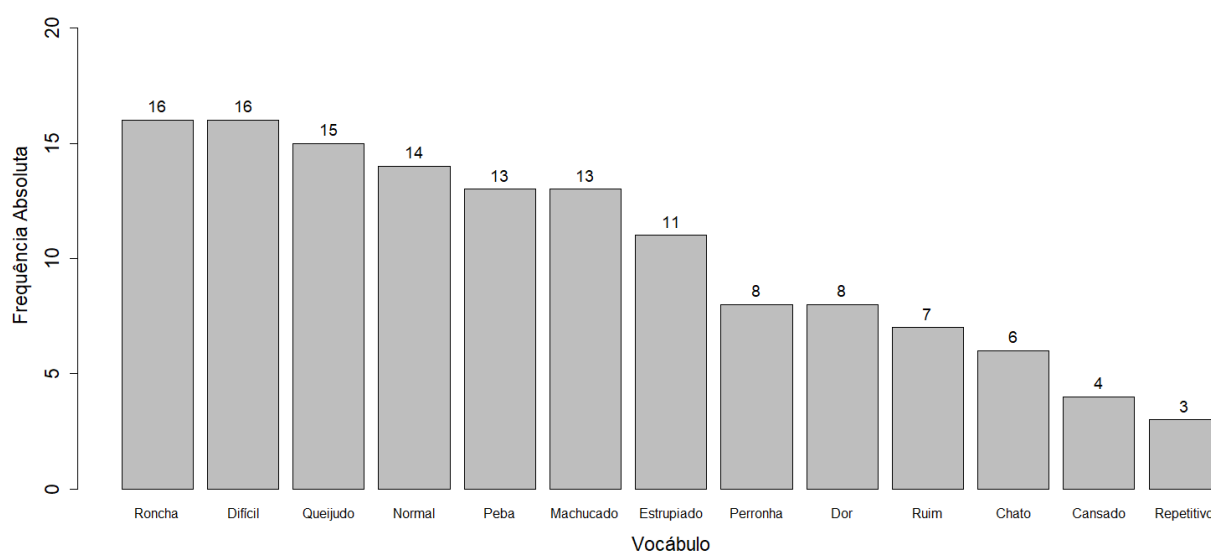


Fonte: Os Autores (2025).

A *tokenização* dos Sentimentos Negativos (Figura 4) revelou possível empate técnico entre as seis primeiras ocorrências, e ratificou a expressão associada às características culturais e sociais pelo uso de vocábulos como *Estrupiado* (machucado), *Peba* (algo ruim, má qualidade), *Perronha* (pessoa ruim no futebol), *Queijudo* (pessoa

fresca) e Roncha (mancha na pele por pancada). Assim como realizado no parágrafo anterior, os destaques em *itálico* não constam no Vocabulário Ortográfico da Língua Portuguesa – VOLP (Academia Brasileira de Letras, 2009), portanto neologismos seriam.

Figura 4: Frequência dos Sentimentos Negativos após Tokenização.



Fonte: Os Autores (2025).

Antes do desenvolvimento do modelo, necessário foi avaliar a distribuição das palavras nos grupos treino, validação e teste, todavia considerando as frequências relativas, dado que os valores absolutos seriam distintos. Explicitando com outras palavras, o percentual de vocábulos em cada conjunto deveria ser similar, o que foi demonstrado pelo teste Qui-quadrado com $\alpha = 0,05$ e 58 graus de liberdade, resultando em valor-p = 0,28. A comparação entre os anos 2023 e 2024, também teve por desfecho a aceitação da hipótese nula (H_0), porque com 29 graus de liberdade, o valor-p = 1,00. Então, o modelo desenvolvido com os dados de 2023 poderia ser aplicado ao banco de dados de 2024.

O conjunto de treino, $n = 75$, foi submetido à validação cruzada *5-fold*, portanto dividido em cinco partes com 15 representantes cada. O modelo conquistou Precisão = $87,87 \pm 2,11\%$, ou seja, dentre as ocorrências previstas como Positivo, a probabilidade de acerto foi de 87,87%, porém dentre todos os casos com similar classificação, o desempenho atenuou, $Recall = 83,64 \pm 2,53\%$, o que era esperado. Isso não comprometeu o equilíbrio do modelo (média harmônica entre Precisão e Recall), pois $F1-Score = 85,70 \pm 2,33\%$. Porém, a capacidade de discriminação entre Positivo e Negativo seria fundamental

em análise de sentimentos, tal característica foi considerada excelente, AUC-ROC = 91,68 $\pm$ 2,33%. Em síntese, os resultados indicaram que o desempenho do modelo deveria ser elevado, dada a ausência de discrepâncias.

Tabela 1: Resultados da Validação Cruzada 5-fold do Conjunto de Treino.

Parte	Precisão	Recall	F1-Score	AUC-ROC
1	85,43	80,72	83,01	89,49
2	89,24	85,29	87,22	92,92
3	85,70	81,04	83,30	89,73
4	89,63	85,76	87,65	93,27
5	89,34	85,41	87,33	93,01
Média	87,87	83,64	85,70	91,68
Desvio Padrão	2,11	2,53	2,33	1,90
Mediana	89,24	85,29	87,22	92,92
Coeficiente de Variação	2,40	3,03	2,72	2,07

Fonte: Os Autores (2025).

O conjunto de validação foi utilizado para ajustar a taxa de aprendizado, hiperparâmetro para ajustar os parâmetros em cada iteração do modelo, ou mais claramente, poderia ser entendida como a rapidez com que o algoritmo atualizaria os próprios parâmetros. Consequentemente, se demasiada anteciparia a convergência, essa se tornaria lenta em taxas baixas (Ottoni *et al.*, 2016). No domínio do corrente estudo, a taxa selecionada foi 0,01 (Tabela 2), porque conquistou o mais alto equilíbrio entre Precisão e Recall. Sequencialmente, a imposição ao conjunto de teste possibilitou a conquista de Precisão = 97,52%, Recall = 92,17%, F1-Score = 94,77 e AUC-ROC = 98,37%.

Tabela 2: Resultados da Taxa de Aprendizado.

Taxa	Precisão	Recall	F1-Score
1.10^{-1}	83,06	78,49	80,71
1.10^{-2}	85,72	82,14	83,89
1.10^{-3}	85,52	78,92	82,09

Fonte: Os Autores (2025).

O modelo foi aplicado ao banco de dados 2024, tendo identificado 84 ocorrências de sentimentos positivos. Posteriormente, profissionais de marketing da escola realizaram a classificação dos 107 registros, determinando 85 como positivos, a divergência se deu sobre uma opinião, a qual considerada foi negativa pelo modelo (Tabela 3). Com isso, a aplicação conquistou Precisão = 100,00%, Recall = 98,82% e F1-Score = 99,41%, ou seja,

a confiabilidade não foi atenuada. A acurácia, total de acertos, considerando todo o banco de dados, foi de, aproximadamente, 99,06%, ratificando a adequação do modelo.

Tabela 3: Resultados Comparação Modelo x Profissional Ano 2024.

	Classes	Profissional		Total
		Negativo	Positivo	
Modelo	Negativo	22	1	23
	Positivo	0	84	84
	Total	22	85	107

Fonte: Os Autores (2024).

Considerações Finais

O modelo de análise de sentimentos desenvolvido pelo método de aprendizagem de máquina se apresentou confiável e detentor de elevada acurácia, portanto com capacidade de escalar as avaliações dos adolescentes da escola esportiva estudada, o que sustentou a conclusão pelo alcance adequado do objetivo.

No contexto de negócio, desenvolver modelos para aspectos específicos do serviço, como satisfação com as intervenções ou treinadores, e qualidade dos equipamentos e instalações pode ser interessante para orientar inovações (disruptivas ou incrementais) e adaptações na prestação do serviço. Investigar com o processamento em linguagem natural a intenção dos responsáveis na permanência ou aquisição do serviço pode instrumentalizar ações de planejamento, comunicação e vendas. Comparar os resultados com modelos desenvolvidos pelos métodos híbrido e baseado em regras permitirá a seleção daquele que forneça melhores resultados, em razão de investimentos em tempo, infraestrutura computacional, capital humano e eficiência de informação.

Referências

ACADEMIA BRASILEIRA DE LETRAS. **VOLP** – Vocabulário Ortográfico da Língua Portuguesa. São Paulo: Globo, 2009.

ADELAKUN, NO. **Métodos de aprendizagem profunda para análise de sentimentos utilizando modelos Bert e Gru**. Edições Nosso Conhecimento, 2024.

AFONSO, AR; DUQUE, CG. Análise de sentimentos em comentários de vídeos do YouTube utilizando aprendizagem de máquina supervisionada. **Ciência e Informação**, v. 48, n. 3, p. 21-33, 2019.

ALCANTARA, HF *et al.* Desempenho em vocabulário receptivo e variáveis sociodemográficas em escolares com queixa de dificuldades de aprendizagem. **Audiology Communication Research**, v. 26, e2523, 2021.

ALMEIDA, PS. Constatações da interferência da oralidade na produção textual através da análise e diagnose de erros. **Revista Iniciação & Formação Docente**, v. 6, n. 1, p. 71-83, 2019.

ALONSO, GI; LETÍCIA, PT. **Análise de sentimento**: um estudo de performance entre comentários e vendas de produtos em um e-Commerce pautado em Text Mining. Trabalho de Conclusão de Curso (Bacharelado em Sistema de Informação) – Faculdade de Computação e Informática. Universidade Presbiteriana Mackenzie. São Paulo, 2022.

ALQURAFI, A; ALSANOOSY, T. Measuring customers' satisfaction using sentiment analysis: model and tool. **Journal of Computer Science**, v. 20, n.4, p. 419-430, 2024.

ALREFAE, R *et al.* Enhancing customer experience through arabic aspect-based sentiment analysis of saudi reviews. **International Journal of Advanced Computer Science & Applications**, v. 15, n. 7, p. 421-427, 2024.

ALVAREZ, DAP; GELBUKH, A; SIDOROV, G. Composer classification using melodic combinatorial n-grams. **Expert Systems with Applications**, v. 249, a. 123300, 2024.

ALVAREZ, KDP. **Uma análise dos sentimentos das empreendedoras de Manaus em busca do equilíbrio trabalho-família**. Trabalho de Conclusão de Curso (Bacharelado em Administração) - Universidade Federal do Amazonas. Manaus (AM), 2024.

ANCHIÊTA, RT; PARDO, TAS. Análise Semântica com base em AMR para o Português. **Linguamática**, v. 14, n. 1, p. 33-48, 2022.

ANTUNES, MBA; ISSA, MF; HOED, RM. Técnicas de machine learning aplicada à mineração de dados e análise de sentimentos para predição de homofobia no twitter. **Revista Foco**, v. 16, n. 1, e853, 2023.

ARÃO, C. Por trás da inteligência artificial: uma análise das bases epistemológicas do aprendizado de máquina. **Trans/formação**, v. 47, n. 3, e02400163, 2024.

ARAÚJO, FLN. **Modelo de arquitetura para uso de banco de dados híbridos adaptável a ferramentas de Processamento de Linguagem Natural (PLN)**. Dissertação (Mestrado Acadêmico em Computação) - Programa de Pós-Graduação em Computação. Universidade Federal do Ceará. Quixadá (CE), 2024.

BATISTA, HR; ARAÚJO, APG. Orações relativas em textos científicos: ambiguidade e efeito discursivo. **Revista do GELNE**, v. 25, n. 3, p. e31377, 2023.

BUDÓ, MLD *et al.* A cultura permeando os sentimentos e as reações frente à dor. **Revista da Escola de Enfermagem da USP**, v. 41, n. 1, p. 36-43, 2007.

BULEGON, H; MORO, CMC. Mineração de texto e o processamento de linguagem natural em sumários de alta hospitalar. **Journal of Health Informatics**, v. 2, n. 2, p. 51-56, 2010.

BUNNELL, P. Dançando com a ambiguidade. **Ambivalências**, v. 12, n. 23, p. 287-303, 2024.

CABRAL, FB. **Interações automatizadas para levantamento de controle de fluxo de processos de negócio**. Dissertação (Mestrado em Computação Aplicada) – Programa de Pós-graduação em Computação Aplicada. Instituto Federal do Espírito Santo. Serra (ES), 2021.

CAMPOS, CCMP. **Avaliação de faqbots através da ferramenta autochatter**. Dissertação (Mestrado em Ciência da Computação) – Pós-graduação em Ciência da Computação. Universidade Federal de Santa Catarina. Florianópolis (SC), 2002.

CAMPOS, EC. **Recomendação de conhecimento disponível em sítios Q&A para auxílio ao desenvolvimento e depuração de software**. Tese (Doutorado em Ciências da Computação) – Programa de Pós-graduação em Ciência da Computação. Faculdade de Computação. Universidade Federal de Uberlândia. Uberlândia (MG), 2019.

CARBONELL, JG; HAYES, PJ. Robust parsing using multiple construction-specific strategies. In BOLC, L (ed.) **Natural language parsing systems**. Berlin (German): Springer Berlin, 1987, p. 1–32.

CONDORI, REL. **Sumarização automática de opiniões baseada em aspectos**. Dissertação (Mestrado em Ciências da Computação e Matemática Computacional) - Instituto de Ciências Matemáticas e de Computação. Universidade de São Paulo. São Carlos (SP), 2015.

CUNHA, JPZ. **Um estudo comparativo das técnicas de validação cruzada aplicadas a modelos mistos**. Dissertação (Mestrado em Ciências) – Instituto de Matemática e Estatística. Universidade de São Paulo. São Paulo, 2019.

DALE, R. Classical approaches to natural language processing. In INDURKHYA, N; DAMERAU, FJ. (Org.). **Handbook of natural language processing**. Florida (USA): Chapman & Hall/CRC, 2010, p. 3-7.

DAVE, K; LAWRENCE, S; PENNOCK, DM. Mining the peanut gallery: opinion extraction and semantic classification of product reviews. In **The Proceedings of the 12th International Conference on World Wide Web**, 2003, pp. 519-528.

DOS SANTOS, CFR. **A robótica educacional como recurso de mobilização e explicitação de invariantes operatórios na resolução de problemas**. Tese (Doutorado em Ensino de Ciência e Tecnologia) Universidade Tecnológica Federal do Paraná. Ponta Grossa (PR), 2018.

EKMAN, P. Facial expression and emotion. **American Psychologist**, v. 48, n. 4, p. 384-392, 1993.

FERNANDES, CM *et al.* **Tipologia de classes e limites da compatibilização das codificações de ocupação nas pesquisas domiciliares do Instituto Brasileiro de Geografia e Estatística (IBGE)**. Brasília (DF): Ipea, 2024. (Diest: Nota Técnica, 64).

FERNANDES, SSZM. **Mineração das publicações relacionadas a Endometriose na rede social Twitter®**. Dissertação (Mestrado em Informática em Saúde) – Programa de Pós-graduação: Mestrado Profissional em Informática em Saúde. Universidade Federal de Santa Catarina. Florianópolis (SC), 2021.

FIORIN, JL. Paixões, afetos, emoções e sentimentos. **CASA: Cadernos de Semiótica Aplicada**, v. 5, n. 2, 2007.

FRANCO, RB; ADANIYA, MHAC. Sistemas de análise de sentimentos usando dados do twitter. **Revista Terra & Cultura: Cadernos de Ensino e Pesquisa**, v. 34, n. especial, p. 111-118, 2013.

FREIRE, RML *et al.* Análise da imagem afetiva dos turistas no destino Vale dos Vinhedos-RS. **Turismo: Visão e Ação**, v. 25, n. 1, p. 114-133, 2023.

GELBUKH, A *et al.* Multi-instrument based N-grams for composer classification task. **Computación y Sistemas**, v. 28, n. 1, p. 85-98, 2024.

GIL, MB. Los límites de la claridad y la comprensión en la documentación administrativa: la notificación de providencia de apremio municipal. **Sphera Publica**, v. 2, n. 24, p. 88-115, 2024.

GOSLING, ITS *et al.* Qualidade em museus: o sentimento dos visitantes do Ecomuseu de Itaipu. **Revista Pretexto**, v. 20, n. 4, p. 89-100, 2019.

HAGHIGHI, A; KLEIN, D. Simple coreference resolution with rich syntactic and semantic features. In **Proceedings of the 2009 Conference on Empirical Methods in Natural Language Processing**, 2009, v. 3, p. 1152–1161.

HIPPISLEY, A. Lexical analysis. In INDURKHYA, N; DAMERAU, FJ. (Org.). **Handbook of natural language processing**. Florida (USA): Chapman & Hall/CRC, 2010, p. 31-58.

INDURKHYA, N; DAMERAU, FJ. (Org.). **Handbook of natural language processing**. Florida (USA): Chapman & Hall/CRC, 2010.

JABEEN, S. Decoding consumer sentiments: Advanced NLP techniques for analyzing smartphone reviews. **Revista de Administração Contemporânea**, v. 28, n. 4, e240102, 2024.

KANG, H; YOO, SJ; HAN, D. Senti-lexicon and improved Naïve Bayes Algorithms for sentiment analysis of restaurant reviews. **Expert Systems with Applications**, v. 39, n. 5, p. 6000–6010, 2012.

KAUTISH, S; KAUR, R. **Análise de sentimentos** - da teoria à prática. Edições Nosso Conhecimento, 2024.

LAPA, ALNS. **Seguindo as máquinas que nos seguem**: considerações sobre a relação entre humanos e não humanos no website Twitter. Trabalho de Conclusão de Curso (Graduação em Ciências Sociais) – Centro de Filosofia e Ciências Humanas. Universidade Federal de Santa Catarina. Florianópolis (SC), 2012.

LAPORTE, É. Resolução de ambiguidades. In RANCHHOD, EM. (Org.) **Tratamento das línguas por computador**. Uma introdução à linguística computacional e suas aplicações. Lisboa (Portugal): Caminho, 2001. p. 49-89.

LAPPIN, S; LEASS, HJ. An algorithm for pronominal anaphora resolution. **Computational Linguistics**, v. 20, n. 4, p. 535–561, 1994.

LIMA, ACES. **Análise de sentimento e desambiguação no contexto da tv social**. Dissertação (Mestrado em Engenharia Elétrica) – Programa de Pós-graduação em Engenharia Elétrica. Universidade Presbiteriana Mackenzie. São Paulo, 2012.

LIU, B. Sentiment analysis and opinion mining. **Synthesis Lectures on Human Language technologies**, v. 5, n. 1, p. 1-167, 2012.

LIU, B. **Sentiment analysis**: Mining opinions, sentiments, and emotions. London (UK): Cambridge University Press, 2015.

LIU, B; ZHANG, L. A survey of opinion mining and sentiment analysis. In AGGARWAL, CC; ZHAI, CX. **Mining text data**. Boston (USA): Springer, 2012, p. 415–463.

LÓPEZ, JO; LÓPEZ, AS; VELÁZQUEZ, RG. **Análise de sentimentos de textos X com aprendizagem profunda**: Um estudo comparativo. Edições Nosso Conhecimento, 2024.

LUDERMIR, TB. Inteligência artificial e aprendizado de máquina: estado atual e tendências. **Estudos Avançados**, v. 35, n. 101, p. 85-94, 2021.

MAIA, NN; SALTON, GD. Utilização de machine learning para classificação de sentimentos no idioma Português-Brasil. **Brazilian Journal of Development**, v. 8, n. 6, p. 43568-43580, 2022.

MALAKIN, LA. **Avaliação de classificadores na análise de sentimentos em redes sociais durante a pandemia da Covid-19**. Dissertação (Mestrado Profissional em Matemática, Estatística e Computação Aplicadas à Indústria) – Instituto de Ciências Matemáticas e de Computação. Universidade de São Paulo. São Carlos (SP), 2024.

MARQUES, BT; MARQUES, LT. **Processamento da linguagem natural**: uma abordagem para avaliação do reconhecimento de entidades nomeadas. Novas Edições Acadêmicas, 2022.

MARTIN, JR; WHITE, PRR. **The language of evaluation**: appraisal in English. New York (USA): Palgrave Macmillan, 2005.

MINKER, W. Stochastic versus rule-based speech understanding for information retrieval. **Speech Communication**, v. 25, n. 4, p. 223 – 247, 1998.

MIONI, JLVM; BARBOSA, CRSC; OLIVEIRA, BFC. Processamento de linguagem natural do português brasileiro para detecção de cyberbullying. **Anais do Computer on the Beach**, v. 15, p. 277-282, 2024.

MORAES, LC *et al.* Análise de ambiguidade linguística em modelos de linguagem de grande escala (LLMs). **Texto Livre**, v. 18, e53181, 2025.

MOURA, RS. **SOS**: Um sistema orto-sintático para a língua portuguesa. Dissertação (Mestrado em Ciência da Computação) - Universidade Federal de Pernambuco. Recife (PE), 1996.

NASCIMENTO, JF; HENZ, RR. A ortografia e os níveis de escrita: o erro em textos de sujeitos escolarizados. **Confluência: Revista do Instituto de Língua Portuguesa**, n. 61, p. 226-248, 2021.

NASUKAWA, T; YI, J. Sentiment analysis: capturing favorability using natural language processing. In **The Proceedings of the 2nd International Conference on Knowledge Capture**, 2003, pp. 70-77.

NEO, G; NEO, A. **Processamento de linguagem natural**: vetorização de texto com Python. Publicação independente, 2024.

NEPOMUCENO, BB *et al.* Bem Estar Pessoal e Sentimento de Comunidade: um estudo psicossocial da pobreza. **Revista Psicologia em Pesquisa**, v. 11, n. 1, p. 74-83, 2017.

OLIVEIRA, DM *et al.* O uso excessivo de tecnologias digitais e seus impactos nas competências socioemocionais de adolescentes. **Revista Contemporânea**, v. 4, n. 12, e6825, 2024.

OTTONI, ALC *et al.* Análise da influência da taxa de aprendizado e do fator de desconto sobre o desempenho dos algoritmos Q-learning e SARSA: aplicação do aprendizado por reforço na navegação autônoma. **Revista Brasileira de Computação Aplicada**, v. 8, n. 2, p. 44-59, 2016.

PALMER, DD. Text preprocessing. In INDURKHYA, N; DAMERAU, FJ. (Org.). **Handbook of natural language processing**. Florida (USA): Chapman & Hall/CRC, 2010, p. 9-30.

PASQUALOTTI, PR; VIEIRA, R. WordnetAffectBR: uma base lexical de palavras de emoções para a língua portuguesa. **Novas Tecnologias na Educação**, v. 6, n. 2, p. 1-10, 2008.

PEDROSA, BAC; DOURADO, JS; LEMOS, SMA. Desenvolvimento lexical, alterações fonoaudiológicas e desempenho escolar: revisão de literatura. **Revista Cefac**, v. 17, n. 5, p. 1633-1642, 2015.

PLUTCHIK, R. The Nature of emotions. **American Scientist**, v. 89, p. 344-350, 2001.

POLO TCF, MIOT HA. Aplicações da curva ROC em estudos clínicos e experimentais. **Jornal Vascular Brasileiro**, v. 19, e20200186, 2020.

QUINTELLA, GM; HANNA, SA. A Lei da Inovação e o Decreto do Marco Regulatório da Inovação: uma análise da jurisprudência do TCU com base em dados proprietária orientada à programação em linguagem nativa. **Cadernos de Prospecção**, v. 17, n. 1, p. 163-175, 2024.

RIBALDO, R; CARDOSO, PCF; PARDO, TAS. Exploring the subtopic-based relationship map strategy for multidocument summarization. **Journal of Theoretical and Applied Computing - RITA**, v. 23, n. 1, p. 183-211, 2016.

RIBEIRO, AP; SILVA, NFF. Um estudo comparativo sobre métodos de análise de sentimentos em tweets. **Revista de Sistema de Informação da FSMA**, n. 22, p. 35-48, 2018.

RIBEIRO, AV *et al.* Classificação de indivíduos ansiosos por árvore de decisão. **Revista Presença**, v. 10, n. 22, p. 64-75, 2024.

RIBEIRO, DS; SÁ, SSL; NOGUEIRA, SM. Ambiguidade de segmentação: casos em que a cadeia falada viabiliza segmentações de sentidos alternativos. **Muiraquitã: Revista de Letras e Humanidades**, v. 12, n. 1, p. 80-92, 2024.

ROAZZI, A *et al.* O que é emoção? Em busca da organização estrutural do conceito de emoção em crianças. **Psicologia: Reflexão e Crítica**, v. 24, p. 51-61, 2011.

ROCHA, M. Relações anafóricas no português falado: uma abordagem baseada em corpus. **DELTA: Documentação de Estudos em Linguística Teórica e Aplicada**, v. 16, n. 2, p. 229-261, 2000.

RODAS, CM *et al.* Análise de sentimentos sobre as vacinas contra covid-19: um estudo com algoritmo de machine learning em postagens no twitter. **Revista de Saúde Digital e Tecnologias Educacionais**, v. 7, n. especial III, p. 24-44, 2022.

ROSA RL; Z RODRÍGUEZ, DZ; BRESSAN, G. **Análise de sentimentos e afetividade nas redes sociais**: métricas de sentimentos e afetividade. Publicação Independente, 2015.

ROSA, JLG. Abordagens ao processamento simbólico da linguagem natural. **Revista do Instituto de Informática da PUC-Campinas**, v. 5, n. 2, p. 20-29, 1997.

SANTOS, IPP. **Análise de sentimento usando redes neurais de convolução**. Dissertação (Mestrado em Engenharia Eletrônica) – Faculdade de Engenharia. Universidade do Estado do Rio de Janeiro. Rio de Janeiro, 2017.

SANTOS, J *et al.* Redes complexas de homônimos para análise semântica textual. **Informação & Informação**, v. 22, n. 1, p. 293-305, 2017.

SILVA, LC *et al.* O impacto das mídias digitais em crianças e adolescentes. **Brazilian Journal of Implantology and Health Sciences**, v. 6, n. 1, p. 1773-1785, 2024.

SILVA, MER; SERRANO, PHSM. Análise de sentimentos em textos de redes sociais: uma comparação entre o ChatGPT e métodos tradicionais. **Cadernos de Comunicação UFSM**, v. 27, n. 3, 2023. doi 10.5902/2316882X84828

SILVA, OB. **Análise de sentimentos**: principais conceitos e aplicação. Monografia (Bacharelado em Engenharia da Computação) – Universidade Federal do Maranhão. São Luís (MA), 2022.

SILVA, TO. **Os impactos sociais, cognitivos e afetivos sobre a geração de adolescentes conectados às tecnologias digitais**. Trabalho de Conclusão de Curso (Graduação em Psicopedagogia) – Universidade Federal da Paraíba. João Pessoa (PB), 2016.

SOARES, KM; TAROUCO, LMR; SILVA, PF. As contribuições de um agente conversacional no ensino e aprendizagem da Física: uma revisão de literatura. **Revista Educar Mais**, v. 5, n. 5, p. 1313-1329, 2021.

SOUZA, AJ. **Comparação entre abordagens de drift detection baseadas em conjuntos de classificadores**: um estudo de caso para previsão de crimes. Dissertação (Mestrado em Informática) – Programa de Pós-graduação em Informática. Pontifícia Universidade Católica do Paraná. Curitiba (PR), 2013.

SOUZA, VA; SOUZA, EF; MEINERZ, GV. Análise de sentimento em tempo real de notícias do mercado de ações. **Brazilian Journal of Development**, v.7, n.1, p. 11084-11091, 2021.

SPROAT, R. **Morphology and computation** (ACL-MIT Series in Natural Language Processing). A Bradford Book, 1992.

STRAUSS, E; JÚNIOR, MVB; FERREIRA, WLL. A importância de utilizar métricas adequadas de avaliação de performance em modelos preditivos de machine learning. **Projectus**, v. 7, n. 2, p. 52-62, 2022.

SUHASINI, V; VIMALA, N. A Hybrid TF-IDF and N-grams based feature extraction approach for accurate detection of fake news on twitter data. **Turkish Journal of Computer and Mathematics Education**, v. 12, n. 6, p. 5710-5723, 2021.

TABOADA, M *et al.* Lexicon-based methods for sentiment analysis. **Computational linguistics**, v. 37, n. 2, p. 267-307, 2011.

TASHTOUSH, Y *et al.* Exploring low-level statistical features of n-grams in phishing URLs: a comparative analysis with high-level features. **Cluster Computing**, v. 27, n. 10, p. 13717-13736, 2024.

THAKUR, S. **Análise de sentimentos sobre as opiniões dos clientes de um sítio de comércio eletrônico**. Limeira (SP): Edições Nosso Conhecimento, 2024.

VICEDO GONZÁLEZ, JL; FERRÁNDEZ RODRÍGUEZ, A; PERAL CORTÉS, J. ¿Cómo influye la resolución de la anáfora pronominal en los sistemas de búsqueda de respuestas?. **Procesamiento del lenguaje natural**, n. 26, p. 231-238, 2000.