

Análise de sentimentos em postagens de redes sociais relacionadas à comunidade LGBTQIA+: desafios, técnicas e contribuições.

Matheus Abreu e Lima de Bairros

Ciência da Computação – Universidade de Passo Fundo (UPF) – Campus
Passo Fundo

172982@upf.br

Abstract. *This thesis examined different sentiment classification methods in social media posts related to the LGBTQIA+ community. The VADER, Random Forest, CNN, and RNN methods were implemented and evaluated. The results highlight the importance of sentiment analysis in this context, with Random Forest achieving the highest accuracy. Sentiment analysis can provide an overview of the emotions and opinions expressed within the LGBTQIA+ community on social media, contributing to understanding their perception and emotional impact. Future research should focus on optimizing the methods and consider ethical aspects to ensure more accurate and inclusive results.*

Resumo. *Este trabalho de conclusão de curso analisou diferentes métodos de análise de sentimentos em postagens de redes sociais relacionadas à comunidade LGBTQIA+. Foram implementados e avaliados os métodos VADER, Random Forest, CNN e RNN. Os resultados destacam a importância da análise de sentimentos nesse contexto e apontam o Random Forest como o método com maior precisão. A análise de sentimentos pode fornecer uma visão geral sobre as emoções e opiniões expressas na comunidade LGBTQIA+ nas redes sociais, contribuindo para compreender sua percepção e impacto emocional. Futuras pesquisas devem focar na otimização dos métodos e considerar aspectos éticos para garantir resultados mais precisos e inclusivos.*

1. Introdução

Com a crescente popularidade das redes sociais, a quantidade de dados gerados diariamente é cada vez mais expressiva. Esses dados são fontes valiosas de informações para empresas, organizações e pesquisadores que desejam entender o comportamento humano e a opinião pública sobre determinados assuntos. No contexto da comunidade LGBTQIA+ (Lésbicas, Gays, Bissexuais, Transgêneros, Queers, Intersexuais, Assexuais), as redes sociais se tornaram um espaço de debate e interação, onde os membros dessa comunidade podem se expressar e compartilhar suas experiências. No entanto, esse espaço também é permeado por discursos de ódio e preconceito, o que pode afetar a saúde mental desses indivíduos.

Nesse sentido, a análise de sentimentos em postagens de redes sociais relacionadas à comunidade LGBTQIA+ é uma área de estudo que ganha cada vez mais importância. Essa técnica permite identificar a polaridade (positiva, negativa ou neutra) dos sentimentos expressos pelos usuários em relação a determinado assunto. Por meio dessa análise, é possível compreender a percepção da sociedade em relação à comunidade LGBTQIA+ e, conseqüentemente, desenvolver ações e políticas públicas

que visem garantir a igualdade de direitos e a proteção dos direitos humanos desses indivíduos.

Neste trabalho de conclusão de curso, serão exploradas técnicas de análise de sentimentos em postagens de redes sociais relacionadas à comunidade LGBTQIA+. Serão apresentados os principais desafios e dificuldades encontrados na aplicação dessas técnicas nesse contexto específico. O objetivo é fornecer uma visão geral sobre a importância da análise de sentimentos em postagens de redes sociais relacionadas à comunidade LGBTQIA+, além de apresentar uma análise crítica das principais técnicas utilizadas e discutir possíveis direções futuras de pesquisa nesse campo de estudo.

2. Contexto Teórico

Nesta sessão, serão abordados os principais conceitos relacionados à análise de sentimentos em postagens de redes sociais relacionadas à comunidade LGBTQIA+. Serão discutidas as técnicas utilizadas para extrair informações sobre a polaridade dos sentimentos expressos pelos usuários, assim como os métodos de aprendizado de máquina aplicados nesse contexto específico.

A análise de sentimentos tem ganhado destaque como uma ferramenta poderosa para compreender a opinião pública e o comportamento humano. Ela permite identificar a polaridade dos sentimentos expressos em textos, como positivos, negativos ou neutros, e fornece insights valiosos sobre a percepção e as emoções dos indivíduos em relação a determinados assuntos.

No contexto das redes sociais, a análise de sentimentos se torna ainda mais relevante, uma vez que essas plataformas se tornaram espaços de interação e debate, onde os membros da comunidade LGBTQIA+ têm a oportunidade de se expressar, compartilhar suas experiências e se conectar com outros indivíduos que compartilham das mesmas vivências.

Entretanto, esses espaços também estão sujeitos à disseminação de discursos de ódio, preconceito e discriminação, o que pode impactar negativamente a saúde mental e o bem-estar emocional dos membros da comunidade. Nesse contexto, a análise de sentimentos em postagens relacionadas à comunidade LGBTQIA+ desempenha um papel importante na identificação e compreensão desses sentimentos, permitindo a implementação de ações e políticas públicas que visem promover a igualdade de direitos e a proteção dos direitos humanos desses indivíduos.

Além disso, este estudo de caso visa explorar as principais técnicas utilizadas na análise de sentimentos, bem como os métodos de aprendizado de máquina aplicados nesse contexto. Serão apresentados os desafios e dificuldades encontrados ao realizar a análise de sentimentos em postagens de redes sociais relacionadas à comunidade LGBTQIA+, assim como as contribuições e resultados obtidos por pesquisas recentes nessa área.

Na próxima sessão, serão discutidas em detalhes as técnicas utilizadas na análise de sentimentos, destacando os métodos de classificação supervisionada, análise léxica e análise de aspectos. Serão apresentados exemplos de algoritmos de aprendizado de

máquina amplamente utilizados nessa área, como Naive Bayes, Máquinas de Vetores de Suporte, Árvores de Decisão e Redes Neurais Artificiais.

2.1 Perplexidade na Análise de Sentimentos

A análise de sentimento, também conhecida como mineração de opinião, é um subcampo da mineração de dados que se concentra na identificação, extração e estudo de informações subjetivas contidas em fontes textuais [Liu 2012]. Essa abordagem tem sido amplamente aplicada em diversas áreas, incluindo marketing, política, saúde e, especialmente, nas redes sociais, com o propósito de compreender a opinião pública acerca de determinados assuntos, identificando tendências e padrões comportamentais.

A perplexidade é um conceito importante na análise de sentimentos, pois mede a incerteza ou surpresa associada a uma sequência de palavras em um texto. Quanto menor a perplexidade, maior a confiança do modelo na classificação do sentimento expresso. A perplexidade é calculada utilizando modelos de linguagem, que atribuem probabilidades a sequências de palavras. Modelos de linguagem baseados em aprendizado de máquina, como os modelos de linguagem neurais, têm se mostrado eficazes na redução da perplexidade e na melhoria do desempenho da análise de sentimentos.

Existem várias abordagens para realizar a análise de sentimentos em textos. Uma das técnicas mais comuns é a classificação supervisionada, em que um modelo de aprendizado de máquina é treinado utilizando um conjunto de dados previamente rotulados, contendo textos classificados de acordo com sua polaridade. O modelo, então, é capaz de generalizar e classificar novos textos com base nos padrões aprendidos durante o treinamento.

Além da classificação supervisionada, outras abordagens incluem a análise lexical, que envolve o uso de léxicos de palavras com polaridade conhecida para atribuir um valor de sentimento a cada palavra no texto, e a análise de aspectos, cujo objetivo é identificar os aspectos específicos mencionados no texto, bem como sua polaridade associada.

2.2 Avanços em Métodos de Aprendizado de Máquina para Análise de Sentimentos

Diversas técnicas de aprendizado de máquina têm sido aplicadas com sucesso na análise de sentimentos. Entre elas, destacam-se algoritmos de classificação, tais como Naive Bayes, Máquinas de Vetores de Suporte (SVM), Árvores de Decisão e Redes Neurais Artificiais.

O algoritmo Naive Bayes é um classificador probabilístico que se baseia no Teorema de Bayes para inferir a classe mais provável de um documento, dadas as palavras presentes nele. Esse algoritmo é reconhecido por sua simplicidade e eficiência computacional, sendo frequentemente utilizado em tarefas de análise de sentimentos.

As SVMs são algoritmos de aprendizado supervisionado que separam os dados em um espaço de alta dimensionalidade, utilizando hiperplanos. Elas são capazes de lidar tanto com problemas de classificação linear quanto não linear, sendo amplamente empregadas em tarefas de análise de sentimentos devido à sua capacidade de lidar com conjuntos de dados de alta dimensionalidade.

As Árvores de Decisão são estruturas de dados que dividem o conjunto de dados em subconjuntos menores, baseando-se em regras de decisão. Essas estruturas são especialmente úteis na análise de sentimentos, uma vez que permitem a interpretação dos resultados de forma intuitiva e transparente.

Redes Neurais Artificiais são modelos inspirados no funcionamento do cérebro humano, compostos por neurônios interconectados. Essas redes têm sido amplamente aplicadas na análise de sentimentos devido à sua capacidade de lidar com problemas complexos e aprender padrões não lineares nos dados. A utilização de redes neurais profundas, como as Redes Neurais Convolucionais (CNNs) e as Redes Neurais Recorrentes (RNNs), tem proporcionado avanços significativos na análise de sentimentos, permitindo uma melhor captura de dependências contextuais e sequenciais nos textos.

2.3 Estudo de Caso: Análise de Sentimentos em Publicações de Redes Sociais Relacionadas à Comunidade LGBTQIA+

No contexto da comunidade LGBTQIA+, as redes sociais desempenham um papel fundamental na expressão e interação de seus membros. Esses espaços online proporcionam uma plataforma para que as pessoas possam compartilhar suas experiências, estabelecer conexões com indivíduos que compartilham vivências semelhantes e também enfrentar desafios, como discursos de ódio e preconceito.

A análise de sentimentos em publicações de redes sociais relacionadas à comunidade LGBTQIA+ tornou-se uma área de estudo relevante, visando compreender a percepção da sociedade em relação a essa comunidade e identificar possíveis impactos na saúde mental dos indivíduos. Através da análise automatizada de sentimentos, é possível extrair informações valiosas sobre o sentimento predominante nas publicações, identificar padrões de discurso e avaliar a polarização e o engajamento em torno de determinados tópicos.

Diversas pesquisas têm abordado essa temática, explorando diferentes técnicas e abordagens de análise de sentimentos. Alguns estudos têm utilizado métodos de aprendizado de máquina, como Naive Bayes, SVM e Redes Neurais, para analisar o sentimento em publicações relacionadas à comunidade LGBTQIA+. Outros estudos têm explorado o uso de recursos léxicos específicos para identificar palavras e termos relacionados a sentimentos positivos ou negativos nas publicações.

Essas pesquisas têm contribuído para uma compreensão mais aprofundada do impacto emocional das interações online relacionadas à comunidade LGBTQIA+. Os resultados obtidos podem ser utilizados para orientar a criação de políticas públicas, campanhas de conscientização e ações de suporte voltadas para essa comunidade,

visando garantir a igualdade de direitos e a promoção do bem-estar emocional dos indivíduos. Além disso, a análise de sentimentos em publicações de redes sociais também pode ajudar a identificar comportamentos abusivos e discursos de ódio, contribuindo para a construção de ambientes online mais inclusivos e seguros.

3. Metodologia

Nesta sessão, serão apresentados os detalhes da metodologia adotada para a realização da análise de sentimentos em postagens de redes sociais relacionadas à comunidade LGBTQIA+. Serão descritos o processo de obtenção e organização do dataset, bem como os métodos de aprendizado de máquina utilizados e a estrutura de implementação empregada.

3.1. Obtenção e Organização do Dataset

O dataset utilizado neste estudo foi obtido a partir do site Kaggle. Inicialmente, o dataset consistia em aproximadamente 1 milhão de postagens na rede social Twitter (tweets) genéricos em inglês, com uma análise preliminar de sentimento dividida em categorias de negativo, positivo, incerteza e litígio. No entanto, para focar na análise de sentimentos relacionada à comunidade LGBTQIA+, foi necessário filtrar os tweets que seriam considerados relevantes no contexto dessa pesquisa, que no caso seriam tweets que façam menção à comunidade LGBTQIA+ em todo e qualquer contexto.

Para realizar a filtragem, foi desenvolvido um código em Python que selecionou apenas os tweets que continham várias palavras-chave relacionadas à comunidade LGBTQIA+. Dessa forma, obteve-se um novo conjunto de dados contendo somente os tweets relevantes para a análise de sentimentos neste estudo.

3.2. Métodos de Aprendizado de Máquina Utilizados

Foram utilizados quatro algoritmos de aprendizado de máquina para realizar a análise de sentimentos nos tweets relacionados à comunidade LGBTQIA+: Valence Aware Dictionary and sEntiment Reasoner (VADER), Random Forest, Rede Neural Convolutacional (CNN) e Rede Neural Recorrente (RNN). A escolha desses algoritmos foi baseada em suas características e capacidades específicas.

O VADER é um modelo baseado em léxico que se mostrou adequado para lidar com textos provenientes de redes sociais, que frequentemente apresentam jargões, emoticons e abreviações [Hutto and Gilbert 2014]. O VADER utiliza um dicionário que atribui pontuações de polaridade a palavras e expressões, permitindo calcular um valor composto de polaridade para cada tweet. Assim, o VADER é capaz de determinar se um tweet expressa um sentimento positivo, negativo ou neutro em relação à comunidade LGBTQIA+.

O algoritmo Random Forest é uma técnica de aprendizado de máquina que cria várias árvores de decisão e faz previsões com base na maioria dos votos dessas árvores [Breiman 2001]. Essa abordagem foi escolhida devido à sua capacidade de lidar com conjuntos de dados volumosos e sua robustez em relação ao sobreajuste. No contexto da análise de sentimentos, o Random Forest foi utilizado para identificar os sentimentos predominantes em relação à comunidade LGBTQIA+ nos tweets.

As RNNs são redes neurais especialmente adequadas para lidar com sequências de dados, como textos, pois são capazes de capturar dependências de longo prazo entre as palavras [Mikolov et al. 2010]. Essa característica as torna uma escolha adequada para a análise de sentimentos em textos provenientes de redes sociais. No contexto deste estudo, as RNNs foram utilizadas para identificar sutilezas e variações nos sentimentos expressos em relação à comunidade LGBTQIA+.

Por fim, as CNNs têm se mostrado eficazes na detecção de características locais em dados e são amplamente utilizadas em tarefas de processamento de linguagem natural [Kim 2014]. A capacidade das CNNs em detectar padrões linguísticos específicos ou expressões emocionais relevantes torna-as uma escolha adequada para a análise de sentimentos. Neste estudo, as CNNs foram aplicadas para extrair informações significativas sobre os sentimentos relacionados à comunidade LGBTQIA+ presentes nos tweets.

4. Implementação dos Algoritmos de Análise de Sentimentos

Neste trabalho de conclusão, foram implementados quatro algoritmos de análise de sentimentos para analisar um conjunto de dados relacionados à comunidade LGBTQIA+ presente em tweets. Os algoritmos utilizados foram o VADER, Random Forest, CNN e RNN.

4.1. Implementação do Algoritmo VADER

O código do algoritmo VADER começa importando as bibliotecas necessárias, como o `nlTK` para a análise de sentimentos baseada em léxico e o `pandas` para manipulação de dados. Também é feito o download do léxico do VADER usando o comando `nlTK.download('vader_lexicon')`.

Em seguida, o conjunto de dados é carregado a partir de um arquivo CSV usando o `pandas`. O conjunto de dados é dividido em conjuntos de treinamento e teste usando o `train_test_split` da biblioteca `scikit-learn`.

O analisador de sentimentos do VADER é inicializado com `SentimentIntensityAnalyzer()`. Em seguida, o algoritmo é aplicado à coluna "Text" dos conjuntos de treinamento e teste. Para cada tweet, são extraídas pontuações de polaridade para as dimensões de sentimento: composto (`compound`), positivo (`positive`), negativo (`negative`) e neutro (`neutral`). Essas pontuações são atribuídas às colunas correspondentes nos conjuntos de dados.

Em seguida, são definidas as características (`features`) e o alvo (`target`) para treinamento do modelo de regressão logística. O modelo de regressão logística é instanciado e treinado usando as características e o alvo do conjunto de treinamento.

Posteriormente, o modelo treinado é usado para fazer previsões no conjunto de teste. A acurácia das previsões é calculada usando a função `accuracy_score` da biblioteca `scikit-learn` e é impressa na saída.

4.2. Implementação do Algoritmo Random Forest

A implementação do algoritmo Random Forest começa importando as bibliotecas necessárias, como o `nlTK` para pré-processamento de texto, o `pandas` para manipulação de dados, o `TfidfVectorizer` para a vetorização de texto, o `train_test_split` da biblioteca

scikit-learn para dividir o conjunto de dados em conjuntos de treinamento e teste, o RandomForestClassifier para a criação do modelo Random Forest e o SMOTE para lidar com o desbalanceamento de dados.

Em seguida, o conjunto de dados é carregado a partir de um arquivo CSV usando o pandas. Também são baixados os recursos necessários do nltk, como stopwords e wordnet, para o pré-processamento do texto.

É definida uma função preprocess_text para realizar o pré-processamento de cada tweet. Nessa função, o texto é convertido para letras minúsculas, a pontuação é removida e as palavras são lematizadas. Além disso, são removidas as stopwords, que são palavras comuns que não agregam muito significado ao sentimento do texto.

O conjunto de dados é então pré-processado aplicando a função preprocess_text à coluna "Text" usando o método apply do pandas. O resultado é armazenado em uma nova coluna chamada "ProcessedText".

O conjunto de dados pré-processado é dividido em conjuntos de treinamento e teste usando o train_test_split da biblioteca scikit-learn. A coluna "ProcessedText" é usada como recurso (feature) e a coluna "Label" é usada como o alvo (target) para a classificação.

Em seguida, é criado um objeto TfidfVectorizer para vetorizar o texto. O vetorizador TF-IDF converte o texto em uma representação numérica, atribuindo pesos a cada palavra com base em sua frequência no documento e em todo o corpus. O vetorizador é ajustado ao conjunto de treinamento e transformado tanto o conjunto de treinamento quanto o conjunto de teste.

Para lidar com o desbalanceamento de dados, é aplicado o Synthetic Minority Over-sampling Technique (SMOTE). O SMOTE gera sinteticamente amostras da classe minoritária, ajudando a equilibrar a distribuição das classes. O conjunto de treinamento é sobre-amostrado usando o SMOTE.

Agora, é instanciado o modelo RandomForestClassifier, que é um conjunto de árvores de decisão aleatórias. O modelo é treinado usando as características (features) e o alvo (target) do conjunto de treinamento.

Em seguida, o modelo treinado é usado para fazer previsões no conjunto de teste. A acurácia das previsões é calculada usando a função accuracy_score da biblioteca scikit-learn e é impressa na saída. Além da acurácia, é gerado um relatório de classificação (classification_report) que inclui métricas como precisão, recall e F1-score para cada classe.

4.3. Implementação do Algoritmo de Rede Neural Convolutacional

A implementação da CNN começa importando as bibliotecas necessárias, como o nltk para o pré-processamento dos textos, pandas para manipulação de dados, e as bibliotecas do Keras para construir e treinar a CNN. Também é feito o download das stopwords e do léxico do WordNet usando os comandos nltk.download('stopwords') e nltk.download('wordnet').

Em seguida, o conjunto de dados é carregado a partir de um arquivo CSV usando o pandas. O pré-processamento dos textos é realizado por meio da remoção de

pontuações, conversão para letras minúsculas e remoção das stopwords. Além disso, é realizada a lematização dos tokens para reduzir as palavras à sua forma básica.

O conjunto de dados é dividido em conjuntos de treinamento e teste usando o `train_test_split` da biblioteca `scikit-learn`.

Em seguida, é definido o número máximo de palavras a serem consideradas (`max_features`), o comprimento máximo de sequência (`maxlen`) e a dimensão de `embedding` (`embedding_dim`) para a camada de `Embedding` da CNN.

É criado um tokenizador do Keras e ajustado aos textos de treinamento. Os textos são convertidos em sequências numéricas com base nas palavras mais frequentes. As sequências numéricas são, então, preenchidas ou truncadas para terem o comprimento máximo especificado.

A classe `SMOTE` do pacote `imblearn` é utilizada para lidar com o desbalanceamento de classes, realizando `oversampling` no conjunto de treinamento.

Em seguida, o modelo da CNN é construído usando a `Sequential API` do Keras. Ele consiste em uma camada de `Embedding`, uma camada `Conv1D` com 128 filtros e função de ativação `ReLU`, uma camada de `pooling` global `MaxPooling1D` e uma camada densa final com ativação sigmoide para a classificação binária.

No modelo da CNN, a camada de `Embedding` é responsável por mapear os índices dos tokens em vetores de `embedding` de tamanho fixo. Essa camada ajuda a representar as palavras em um espaço vetorial contínuo, capturando relações semânticas e semelhanças entre elas. O tamanho dos vetores de `embedding` é definido como parte da configuração do modelo e pode ser ajustado conforme necessário.

A camada `Conv1D` (convolução unidimensional) é uma camada de convolução que atua em sequências de dados. Nesse contexto, ela é usada para extrair características relevantes dos textos, capturando padrões locais em diferentes partes das sequências. A camada `Conv1D` usa filtros para realizar convolução em janelas deslizantes ao longo das sequências, produzindo um mapa de características que destaca informações importantes.

No caso específico mencionado, o modelo da CNN utiliza 128 filtros na camada `Conv1D`. Esses filtros são responsáveis por extrair diferentes características dos textos analisados. Cada filtro captura informações específicas, como palavras-chave, estruturas gramaticais ou outros padrões relevantes para a classificação de sentimentos. A utilização de múltiplos filtros permite que o modelo aprenda uma variedade de características em paralelo, aumentando a capacidade de generalização e melhorando a precisão da classificação.

Os rótulos de classe são convertidos para valores numéricos: 0 para "negative", 1 para "positive" e 0.5 para "uncertainty" e "litigious".

O modelo é compilado com o otimizador Adam, uma taxa de aprendizado de 0.001 e a função de perda `binary_crossentropy`. Esses parâmetros são selecionados para otimizar o processo de treinamento do modelo. A taxa de aprendizado determina o tamanho dos passos dados durante o treinamento, influenciando na velocidade e na estabilidade da convergência do modelo. A taxa de aprendizado de 0.001 é escolhida para controlar cuidadosamente o ajuste dos pesos do modelo.

Quanto às épocas e ao tamanho do lote, eles são definidos como 10 e 32, respectivamente. As épocas representam o número de vezes que o conjunto de treinamento é percorrido durante o treinamento do modelo. Ter 10 épocas permite que o modelo explore melhor os padrões dos dados e refine suas habilidades de classificação. Já o tamanho do lote determina o número de amostras de treinamento processadas antes de atualizar os pesos do modelo. O tamanho de lote de 32 é escolhido para equilibrar a eficiência computacional e a estabilidade do treinamento, permitindo atualizações frequentes dos pesos com base em um número razoável de amostras. Durante o treinamento, a acurácia do modelo também é calculada para avaliar o desempenho da classificação.

Após o treinamento, o modelo é usado para fazer previsões no conjunto de teste. As previsões são convertidas para valores binários (0 ou 1) com base em um limite de 0.5. Os rótulos originais são convertidos de volta para valores binários (0 ou 1), substituindo o valor 0.5 por 0. A acurácia das previsões é calculada usando a função `accuracy_score` da biblioteca `scikit-learn` e é impressa na saída. Além disso, é gerado um relatório de classificação usando a função `classification_report` da biblioteca `scikit-learn`, que exibe métricas como precisão, recall e F1-score para cada classe.

4.4. Implementação do Algoritmo de Rede Neural Recorrente

O código do algoritmo RNN começa importando as bibliotecas necessárias, como o `nlTK` para a análise de sentimentos baseada em léxico e o `pandas` para manipulação de dados. Também é feito o download dos stopwords e do `WordNetLemmatizer` do NLTK usando o comando `nlTK.download('stopwords')` e `nlTK.download('wordnet')`.

Em seguida, o conjunto de dados é carregado a partir de um arquivo CSV usando o `pandas`. O conjunto de dados é dividido em conjuntos de treinamento e teste usando o `train_test_split` da biblioteca `scikit-learn`.

É definido o número máximo de recursos (`max_features`), o tamanho máximo das sequências (`maxlen`) e a dimensão do embedding (`embedding_dim`) para o modelo RNN.

Em seguida, é criado um objeto `Tokenizer` para converter os textos em sequências de tokens e atribuir um índice a cada palavra. O `Tokenizer` é ajustado aos textos do conjunto de treinamento e, em seguida, os textos do conjunto de treinamento e teste são convertidos em sequências de tokens usando o método `texts_to_sequences`.

As sequências de tokens são então preenchidas para terem o mesmo comprimento (`maxlen`) usando o método `pad_sequences`. Isso garante que todas as sequências tenham o mesmo tamanho, permitindo que sejam usadas como entrada para o modelo RNN.

Em seguida, o balanceamento de dados é tratado usando o método `SMOTE` da biblioteca `imbalanced-learn`. O `SMOTE` é aplicado ao conjunto de treinamento, gerando novas amostras sintéticas para a classe minoritária, a fim de equilibrar as classes.

O modelo RNN é criado usando a classe `Sequential` do `Keras`. É adicionada uma camada de embedding que mapeia os índices dos tokens para vetores de embedding de tamanho `embedding_dim`. Em seguida, é adicionada uma camada `LSTM` (Long

Short-Term Memory) com 128 unidades. A camada LSTM é uma camada de células de memória recorrentes que é capaz de lembrar informações de longo prazo. Por fim, é adicionada uma camada densa com ativação sigmoid para a classificação binária.

Os rótulos de classe (`y_train` e `y_test`) são convertidos em valores numéricos, em que a classe "negative" é mapeada para 0, a classe "uncertainty" e "litigious" são mapeadas para 0.5 e a classe "positive" é mapeada para 1.

O modelo é compilado com o otimizador Adam e a função de perda `binary_crossentropy`. Em seguida, o modelo é treinado usando os dados de treinamento. São especificadas 10 épocas de treinamento e um tamanho de lote de 32. Também é definida uma validação de 20% dos dados de treinamento para monitorar o desempenho do modelo durante o treinamento.

Após o treinamento, assim como no algoritmo da CNN, o modelo é usado para fazer previsões no conjunto de teste. As previsões são convertidas para valores binários (0 ou 1) com base em um limite de 0.5. Os rótulos originais são convertidos de volta para valores binários (0 ou 1), substituindo o valor 0.5 por 0. A acurácia das previsões é calculada usando a função `accuracy_score` da biblioteca `scikit-learn` e é impressa na saída. Além disso, é gerado um relatório de classificação usando a função `classification_report` da biblioteca `scikit-learn`, que exibe métricas como precisão, recall e F1-score para cada classe.

5. Resultados e Discussões

Com base nos quatro métodos implementados, VADER, Random Forest, CNN e RNN, foram obtidos resultados diferentes para cada um deles. Para uma melhor compreensão dos resultados, é importante entender as métricas de avaliação utilizadas na análise desses métodos.

A precisão (`accuracy`) é uma métrica que mede a proporção de amostras corretamente classificadas em relação ao total de amostras. Ela fornece uma visão geral do desempenho geral do modelo. Uma alta precisão indica um bom desempenho na classificação dos sentimentos.

O relatório de classificação (`classification report`) é uma análise mais detalhada que apresenta métricas como precisão (`precision`), recall (`revocação`) e F1-score para cada classe. A precisão mede a proporção de verdadeiros positivos em relação ao total de classificações positivas feitas pelo modelo. Ela indica a capacidade do modelo de classificar corretamente uma determinada classe. O recall mede a proporção de verdadeiros positivos em relação ao total de amostras que pertencem a uma classe específica. Ele indica a capacidade do modelo de recuperar corretamente amostras de uma classe. O F1-score é uma medida de equilíbrio entre precisão e recall, fornecendo uma avaliação geral do desempenho do modelo.

Além disso, o relatório de classificação também apresenta a coluna de suporte (`support`), que indica a contagem de amostras de cada classe no conjunto de teste. Isso pode ser útil para identificar classes com desequilíbrio de dados.

Ao comparar os resultados dos métodos implementados, é possível observar o desempenho de cada um em termos de precisão, recall e F1-score para cada classe.

Essas métricas permitem uma análise mais aprofundada do desempenho dos modelos em diferentes aspectos.

É importante destacar que essas métricas são apenas algumas das muitas que podem ser utilizadas na análise de resultados. Outras métricas, como acurácia ponderada (weighted accuracy), macro média (macro average) e média ponderada (weighted average), também podem ser relevantes dependendo do contexto do estudo.

Com base na análise detalhada dos resultados obtidos, é possível obter insights sobre a eficácia de cada método na classificação de sentimentos. Essas informações podem ser úteis para compreender a capacidade dos modelos implementados em lidar com as particularidades dos dados e melhorar sua precisão e desempenho.

No entanto, é importante ressaltar que a análise de resultados não se limita apenas às métricas mencionadas. É necessário considerar outros fatores, como o tamanho e a qualidade do conjunto de dados, a escolha adequada dos parâmetros e hiperparâmetros dos modelos, bem como a interpretação dos resultados à luz do contexto específico do estudo.

Dessa forma, a análise dos resultados é uma etapa crucial na avaliação e compreensão dos métodos implementados, fornecendo insights valiosos sobre seu desempenho e direcionando possíveis melhorias e ajustes para pesquisas futuras.

5.1. Método VADER

O método VADER é uma ferramenta de análise de sentimentos que atribui pontuações de polaridade aos textos. A precisão alcançada pelo modelo implementado utilizando o VADER foi de aproximadamente 0,4956.

5.2. Método Random Forest

O método Random Forest é um algoritmo de aprendizado de máquina baseado em árvores de decisão. Ele foi utilizado neste trabalho de conclusão para realizar a classificação de sentimentos. A precisão obtida com o modelo de Random Forest foi de aproximadamente 0,9443. Além disso, foi gerado um relatório de classificação que apresentou métricas como precisão, recall e F1-score para cada classe. Observou-se um desempenho satisfatório para todas as classes, com valores acima de 0,9 na maioria das métricas.

Tabela 1. Relatório de classificação gerado pelo método Random Forest

	Precision	Recall	F1-Score	Support
Litigious	0.93	0.95	0.94	98
Negative	0.94	0.95	0.94	205
Positive	0.98	0.94	0.96	141
Uncertainty	0.92	0.94	0.93	131
Accuracy			0.94	575

Macro Avg.	0.94	0.94	0.94	575
Weighted Avg.	0.94	0.94	0.94	575

5.3. Método Rede Neural Convolucional

A CNN é um tipo de arquitetura de rede neural especialmente adequada para processamento de dados em forma de grade, como imagens ou sequências de palavras. Neste trabalho, ela foi implementada para a classificação de sentimentos. A precisão alcançada com o modelo de CNN foi de aproximadamente 0,8713. O relatório de classificação gerado mostrou um desempenho satisfatório para a classe negativa, com uma alta precisão. No entanto, a classe positiva teve um desempenho um pouco inferior, com uma precisão mais baixa.

Tabela 2. Relatório de classificação gerado pelo método CNN

	Precision	Recall	F1-Score	Support
0.0	0.99	0.84	0.91	333
1.0	0.66	0.97	0.79	110
Accuracy			0.87	443
Macro Avg.	0.83	0.91	0.85	443
Weighted Avg.	0.91	0.87	0.88	443

5.4. Método Rede Neural Recorrente

A RNN é um tipo de arquitetura de rede neural projetada para processar dados sequenciais. Neste trabalho, a RNN foi implementada para a classificação de sentimentos. A precisão obtida com o modelo de RNN foi de aproximadamente 0,8329. O relatório de classificação gerado revelou um desempenho satisfatório para a classe negativa, com uma alta precisão. No entanto, assim como na CNN, a classe positiva apresentou um desempenho inferior, com uma precisão mais baixa.

Tabela 3. Relatório de classificação gerado pelo método RNN

	Precision	Recall	F1-Score	Support
0.0	0.95	0.82	0.88	333
1.0	0.62	0.87	0.72	110
Accuracy			0.83	443
Macro Avg.	0.78	0.85	0.80	443

Weighted Avg.	0.87	0.83	0.84	443
---------------	------	------	------	-----

5.5. Análise Comparativa

Ao comparar os resultados dos métodos implementados, observamos que o Random Forest obteve a maior precisão entre todos os métodos, com aproximadamente 0,9443. Isso indica que o modelo de Random Forest teve um desempenho melhor na classificação dos sentimentos em comparação com os outros métodos.

A CNN e a RNN obtiveram resultados de precisão um pouco inferiores, com aproximadamente 0,8713 e 0,8329, respectivamente. No entanto, é importante notar que a CNN teve um desempenho melhor na classe negativa, enquanto a RNN teve um desempenho um pouco melhor na classe positiva.

O método VADER alcançou uma precisão de 0,4957, sendo o método que obteve a menor precisão entre todos os modelos testados. Isso indica que o modelo teve dificuldade em classificar corretamente as amostras do conjunto de teste.

Em termos gerais, os resultados indicam que os modelos implementados, com exceção do VADER, têm a capacidade de realizar a classificação de sentimentos, mas com variações no desempenho para cada método. O Random Forest se destacou como o método com a maior precisão, enquanto a CNN e a RNN apresentaram desempenho inferiores, porém ainda satisfatórios.

6. Conclusão

Este artigo analisou diferentes métodos de classificação de sentimentos em postagens de redes sociais relacionadas à comunidade LGBTQIA+. Os resultados obtidos fornecem uma visão geral sobre a importância da análise de sentimentos nesse contexto e apresentam uma análise crítica das principais técnicas utilizadas.

Os métodos avaliados incluíram o VADER, Random Forest, CNN e RNN. Embora os resultados específicos tenham sido detalhados nas seções anteriores, é importante ressaltar que o Random Forest obteve a maior precisão, seguido pela CNN e RNN. O método VADER apresentou a menor precisão.

Esses resultados demonstram a relevância da análise de sentimentos para compreender as emoções e opiniões expressas nas postagens relacionadas à comunidade LGBTQIA+ nas redes sociais. A capacidade de identificar sentimentos positivos e negativos pode ser fundamental para avaliar o impacto emocional dessas postagens e entender a percepção da comunidade em relação a diferentes questões.

Embora as técnicas avaliadas tenham mostrado resultados promissores, é importante destacar que ainda há espaço para melhorias e avanços. A seleção cuidadosa do método de análise de sentimentos, considerando as peculiaridades do conjunto de dados e as necessidades específicas, pode contribuir para resultados mais precisos e relevantes.

Além disso, existem várias direções futuras de pesquisa nesse campo. A otimização das arquiteturas de rede neural, o refinamento dos algoritmos de aprendizado de máquina e a incorporação de abordagens de processamento de linguagem natural

mais avançadas podem aprimorar ainda mais a precisão e eficiência dos modelos de análise de sentimentos.

Também é importante considerar aspectos éticos, como a garantia da imparcialidade e equidade nas análises de sentimentos, evitando vieses indesejados. A inclusão de mais dados diversificados e representativos da comunidade LGBTQIA+ pode ajudar a abordar essas questões e fornecer resultados mais justos e inclusivos.

Em suma, os resultados deste estudo destacam a importância da análise de sentimentos em postagens relacionadas à comunidade LGBTQIA+ nas redes sociais. Essa análise pode fornecer insights valiosos para pesquisadores, ativistas e profissionais que desejam compreender melhor as experiências e perspectivas dessa comunidade. Com o contínuo avanço da tecnologia e a pesquisa nesse campo, espera-se que a análise de sentimentos se torne uma ferramenta cada vez mais poderosa para promover a inclusão e a igualdade.

Referências

- Liu, B. “Sentiment Analysis: A Fascinating Problem”, In: Sentiment Analysis and Opinion Mining, Morgan & Claypool Publishers, 2012.
- Hutto, C. and Gilbert, E. “VADER: A Parsimonious Rule-Based Model for Sentiment Analysis of Social Media Text”, Proceedings of the International AAAI Conference on Web and Social Media, p. 216-225, 2014.
- Breiman, L. “Random Forests”, Machine Learning 45, p. 5–32, 2001.
- Mikolov, T., Karafiat, M., Burget, L., Cernocký, J. and Khudanpur, S. “Recurrent neural network based language model”, Interspeech Vol. 2. No. 3., p. 1045-1048, 2010.
- Kim, Y. “Convolutional Neural Networks for Sentence Classification”, Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP), p. 1746–1751, 2014.