



**INSTITUTO
FEDERAL**

Piauí

**INSTITUTO FEDERAL DE EDUCAÇÃO, CIÊNCIA E TECNOLOGIA DO PIAUÍ
CAMPUS PICOS
ANÁLISE E DESENVOLVIMENTO DE SISTEMAS**

ÊMYLLE BEATRIZ DE SOUSA

**CLASSIFICAÇÃO DE TROCADILHOS EM PORTUGUÊS COM
BERTIMBAU LARGE: DESAFIOS E RESULTADOS**

**PICOS, PIAUÍ
2025**

ÊMYLLE BEATRIZ DE SOUSA

CLASSIFICAÇÃO DE TROCADILHOS EM PORTUGUÊS COM
BERTIMBAU LARGE: DESAFIOS E RESULTADOS

Trabalho de Conclusão de Curso (Artigo Científico) apresentado como exigência parcial para obtenção do diploma do Curso de Análise e Desenvolvimento de Sistemas do Instituto Federal de Educação, Ciência e Tecnologia do Piauí, Campus Picos.

Orientador: Prof. M^e. Aislan Rafael Rodrigues Sousa

Coorientador: Prof. Dr. Rafael Torres Anchiêta

PICOS, PIAUÍ
2025

ÊMYLLE BEATRIZ DE SOUSA

CLASSIFICAÇÃO DE TROCADILHOS EM PORTUGUÊS COM
BERTIMBAU LARGE: DESAFIOS E RESULTADOS

Trabalho de Conclusão de Curso (Artigo Científico) apresentado como exigência parcial para obtenção do diploma do Curso de Análise e Desenvolvimento de Sistemas do Instituto Federal de Educação, Ciência e Tecnologia do Piauí, Campus Picos.

Aprovado em 02 de Setembro de 2025.

BANCA EXAMINADORA:

Prof. M^e. Aislan Rafael Rodrigues Sousa(Orientador)
Instituto Federal do Piauí (IFPI)

Prof. Dr. Rafael Torres Anchiêta (Co-orientador)
Instituto Federal do Maranhão (IFMA)

Prof. M^e. Jáder Anderson Oliveira de Abreu
Instituto Federal do Piauí (IFPI)

Prof. M^e. Joao Paulo Lima do Nascimento
Instituto Federal do Piauí (IFPI)

PICOS, PIAUÍ
2025

Dedico este trabalho à versão de mim que ele ajudou a construir. Que esta conquista não seja um ponto de chegada, mas o alicerce que te impulsionará a voos ainda mais altos. E que, ao colher os frutos, você se lembre com orgulho da semeadura.

Honre a jornada.

CLASSIFICAÇÃO DE TROCADILHOS EM PORTUGUÊS COM BERTIMBAU LARGE: DESAFIOS E RESULTADOS

Émylle Beatriz de Sousa

Prof. M^º. Aislan Rafael Rodrigues Sousa | Prof. Dr. Rafael Torres Anchiêta

RESUMO

Este trabalho apresenta um estudo sobre a detecção automática de trocadilhos em português, um desafio significativo no campo do Processamento de Linguagem Natural (PLN) devido à ambiguidade semântica e à dependência de contexto. Para a análise, foi utilizado o corpus Puntuguese, uma coleção inédita composta por 3.990 textos curtos, igualmente divididos entre trocadilhos e não trocadilhos, provenientes de diferentes fontes como blogs, redes sociais e vídeos no YouTube. O modelo BERTimbau Large, especializado nas variações do português brasileiro e europeu, foi aplicado para analisar as sentenças. O trabalho também envolveu técnicas de pré-processamento e a extração de padrões linguísticos por expressões regulares, que ajudaram na análise dos dados. A tarefa de identificar trocadilhos é desafiadora devido às particularidades e sutilezas que esse tipo de humor envolve, exigindo uma compreensão profunda das estruturas linguísticas. Este estudo oferece uma contribuição importante para o avanço da detecção de humor verbal em português e abre portas para futuros desenvolvimentos na área.

Palavras-chaves: processamento de linguagem natural (PLN); detecção automática de trocadilhos; bertimbau large; humor verbal computacional; classificação de textos em português.

ABSTRACT

This paper presents a study on the automatic detection of puns in Portuguese, a significant challenge in the field of Natural Language Processing (NLP) due to semantic ambiguity and context dependence. For the analysis, we used the Puntuguese corpus, a unique collection of 3,990 short texts, equally divided between puns and non-puns, from various sources such as blogs, social media, and YouTube videos. The BERTimbau Large model, specialized in variations of Brazilian and European Portuguese, was applied to analyze the sentences. The work also involved preprocessing techniques and the extraction of linguistic patterns using regular expressions, which aided in data analysis. The task of identifying puns is challenging due to the particularities and subtleties of this type of humor, requiring a deep understanding of linguistic structures. This study offers an important contribution to the advancement of verbal humor detection in Portuguese and opens doors for future developments in the field.

Keywords: natural language processing (NLP); automatic detection of puns; large bertimbau; computational verbal humor; text classification in Portuguese.

Data de Aprovação: 02 de Setembro de 2025

1 INTRODUÇÃO

O humor, por sua natureza ambígua e subjetiva, representa um dos maiores desafios para o Processamento de Linguagem Natural (PLN). Entre suas manifestações, os trocadilhos se destacam como uma forma de jogo de palavras (*wordplay*) que exige dos sistemas computacionais uma compreensão profunda das estruturas linguísticas e dos múltiplos sentidos que elas podem carregar. Como define Delabastita (2016):

“jogo de palavras é o nome geral para os diversos fenômenos textuais em que características estruturais da(s) língua(s) utilizada(s) são exploradas com o objetivo de provocar uma confrontação comunicativamente significativa entre duas (ou mais) estruturas linguísticas com formas mais ou menos semelhantes e significados mais ou menos diferentes”.

Essa confrontação entre forma e significado é o que torna os trocadilhos particularmente desafiadores para sistemas de PLN, pois exige não apenas reconhecimento fonológico, mas também interpretação semântica contextualizada.

A motivação para a detecção automática de trocadilhos se fundamenta na crescente demanda por sistemas capazes de compreender e interagir de maneira mais natural com o conteúdo textual gerado em redes sociais, memes, *chats* e outras plataformas digitais. Em diversas aplicações, como moderação de conteúdo, geração de texto criativo, assistentes virtuais e análise de sentimentos, a identificação precisa de elementos humorísticos — como os trocadilhos — pode melhorar significativamente a qualidade da interação e a precisão do sistema.

No entanto, o reconhecimento de trocadilhos apresenta desafios consideráveis: sua natureza pode ser sutil, muitas vezes dependendo de duplo sentido ou conhecimento prévio, e as variações regionais e contextuais na língua portuguesa podem tornar o processo ainda mais complexo.

Para resolver esse problema, este trabalho apresenta um sistema de classificação de textos voltado para a detecção automática de trocadilhos, utilizando o modelo de linguagem *BERTimbau Large* Souza, Nogueira & Lotufo (2020). O *BERTimbau Large* é uma adaptação do BERT, treinado especificamente para o português, que se destaca em tarefas de PLN, como compreensão contextual e reconhecimento de padrões semânticos em sentenças. Esse modelo, treinado no *corpus Puntuguese* Inacio *et al.* (2024), especializado em português, tem se mostrado eficiente em várias tarefas de PLN, como a identificação de discurso de ódio, similaridade textual e reconhecimento de entidades nomeadas. Além disso, o modelo permite a captura de padrões linguísticos específicos, como jogos de palavras e ambiguidades, que não são detectáveis apenas pela tokenização. Essas técnicas ajudaram a tornar o sistema mais robusto e preciso, permitindo identificar com maior acuracidade os trocadilhos presentes nas sentenças.

Embora a detecção de trocadilhos seja uma tarefa desafiadora, este estudo se baseia no trabalho de Inacio *et al.* (2024), que utilizou o *corpus Puntuguese*, mas apresenta avanços significativos. O modelo proposto superou os resultados anteriores

ao identificar com maior precisão os trocadilhos presentes nas sentenças. Esses resultados preliminares indicam que, embora haja espaço para melhorias, o estudo representa um passo importante na evolução das técnicas de detecção de humor verbal, especialmente no contexto da língua portuguesa. Este trabalho também destaca a complexidade do *corpus* e as oportunidades de aprimoramento nas abordagens atuais.

A estrutura do artigo está organizada da seguinte forma: no Capítulo 2, são revisados os principais trabalhos relacionados à detecção de humor e trocadilhos; no Capítulo 3, é descrito em detalhes o corpus Puntuguese, suas características e origens; no Capítulo 4, são apresentadas a metodologia adotada e os ajustes realizados no modelo BERTimbau Large; no Capítulo 5, são discutidos os resultados experimentais, suas implicações e limitações; por fim, no Capítulo 6, são apresentadas as conclusões e apontadas direções para futuras pesquisas.

2 TRABALHOS RELACIONADOS

A detecção de humor e, mais especificamente, a identificação de trocadilhos em textos é uma tarefa desafiadora no campo do Processamento de Linguagem Natural (PLN). A ambiguidade linguística, a dependência de contextos culturais e a natureza subjetiva do humor tornam essa tarefa particularmente difícil para os sistemas automáticos. Diversas abordagens têm sido exploradas para lidar com esses desafios, com modelos baseados em aprendizado de máquina e redes neurais se destacando como alternativas eficazes.

O trabalho propõe uma competição focada na detecção e interpretação de trocadilhos em inglês. Este estudo representa um avanço significativo na área, fornecendo uma avaliação ampla sobre diferentes abordagens para detectar trocadilhos e destacando a importância do contexto e das características semânticas para essa tarefa Miller, Hempelmann & Gurevych (2017).

Embora o estudo se concentre no inglês, ele fornece metodologias e desafios valiosos que podem ser adaptados à detecção de trocadilhos em outros idiomas, como o português, embora seja necessário ajustar as abordagens para as particularidades linguísticas e culturais de cada língua.

Em uma linha de pesquisa semelhante sobre a detecção automática de humor, Bonet, Rincón e López (2023) investigaram a identificação de humor prejudicial em textos do Twitter, utilizando dados em espanhol. Embora o estudo não se concentre em trocadilhos, ele aborda desafios similares relacionados às nuances semânticas e contextuais do humor textual Bonet, Rincón & López (2023). Os autores obtiveram resultados expressivos, medidos por uma métrica que combina precisão e recall de todas as classes de forma equilibrada (*F1-Macro*), atingindo 89,5%, evidenciando o potencial dos modelos baseados em transformadores. Ainda assim, o trabalho aponta

limitações relacionadas ao ajuste excessivo aos dados de treinamento — conhecido como *overfitting* — especialmente quando o conjunto de dados é pequeno, um problema também recorrente na detecção de trocadilhos.

Por outro lado, Cruz *et al.* (2023) exploraram o uso de *ensembles*, que combinam múltiplos modelos para melhorar a robustez das previsões, aplicados a *tweets* em espanhol no contexto da competição HUH@IberLEF 2023. A abordagem utilizou modelos como *BERT*, um modelo pré-treinado baseado em redes do tipo transformador capaz de capturar o contexto de cada palavra considerando tanto termos anteriores quanto posteriores na frase, além de *RoBERTa* e *BETO*.

O sistema alcançou *F1-Macro* de 79,6% na tarefa de classificação multilabel dos grupos-alvo, obtendo o 1º lugar entre 49 equipes participantes. Essa estratégia de combinar múltiplos modelos mostra-se promissora para a detecção de trocadilhos, pois permite capturar diferentes nuances linguísticas e contextuais, essenciais nesse tipo de humor.

O Puntuguese Inacio *et al.* (2024) se destaca como um avanço relevante na detecção de trocadilhos em português. Desenvolvido para resolver limitações de corpora anteriores, como o utilizado por Gonçalo Oliveira *et al.* (2020), o Puntuguese passou por curadoria manual para evitar que os modelos aprendessem padrões superficiais, o que prejudicava sua capacidade de generalização — ou seja, de acertar exemplos novos de forma consistente.

Com textos extraídos de fontes populares — *blogs*, Instagram e YouTube — e organização balanceada entre exemplos com e sem trocadilhos, o corpus se consolidou como uma base robusta e amplamente utilizada em estudos sobre humor verbal no contexto do PLN em português.

Este artigo se diferencia principalmente por seu foco exclusivo na detecção de trocadilhos em português, utilizando o corpus Puntuguese, que oferece um conjunto de dados balanceado e representativo para a tarefa. Enquanto muitos estudos anteriores tratam do humor de maneira geral, a presente pesquisa se dedica a compreender as sutilezas semânticas dos trocadilhos, um tipo de humor que exige uma interpretação mais profunda.

Para isso, adotou-se o modelo BERTimbau large, que foi treinado especificamente para o português, e foi aplicado junto com técnicas de extração de características linguísticas por expressões regulares, permitindo uma análise mais detalhada e precisa dos textos.

3 DESCRIÇÃO DO CONJUNTO DE DADOS

Para a realização deste trabalho, foi utilizado o corpus Puntuguese (INACIO *et al.*, 2024), uma coleção única de textos de trocadilhos em português, abrangendo tanto o português brasileiro quanto o europeu. Este corpus foi construído a partir de três fontes principais: o *blog* Maiores e Melhores, a página do Instagram O Sagrado Caderno das

Piadas Secas, e o canal *UTC–Ultimate Trocadilho Challenge* dos Castro Brothers no YouTube. O corpus contém 4.903 textos, classificados em duas categorias: trocadilho, com 2.450 amostras, e não trocadilho, com 2.453 amostras, o que proporciona um corpus balanceado, sem a necessidade de técnicas adicionais de balanceamento de classes. A seguir, seguem exemplos representativos de cada classe do corpus:

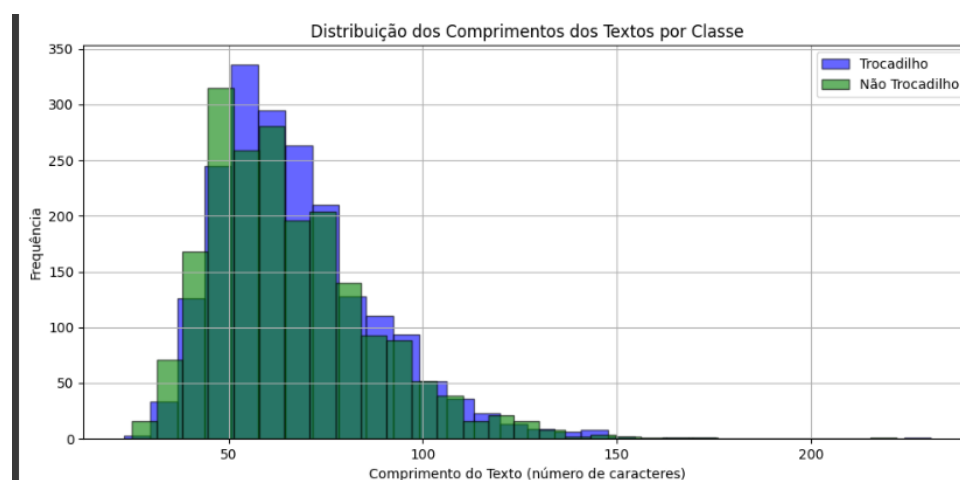
- **Exemplo de trocadilho:** “Por que a abelha foi ao médico? Porque estava com um zangão na garganta.”
- **Exemplo de não trocadilho:** “Ela foi ao cinema ver um filme de drama com as amigas.”

Esses exemplos ilustram a diferença entre construções textuais com jogos de palavras — que exploram ambiguidade fonológica ou semântica — e frases descritivas comuns, sem intenção humorística. Observa-se que a presença de perguntas estruturadas, como no primeiro exemplo, é frequente entre os trocadilhos, o que pode ter influenciado o modelo a superestimar o humor em determinados contextos.

As sentenças do corpus são predominantemente curtas, como evidenciado pela análise de comprimento de texto. A maioria das sentenças possui entre 50 e 100 caracteres, e a quantidade de sentenças diminui à medida que o comprimento das mesmas aumenta. Apenas um pequeno número de textos ultrapassa os 150 caracteres, indicando que os trocadilhos são, em sua maioria, textos curtos ou de tamanho médio.

A Figura 1 apresenta a distribuição dos comprimentos dos textos, em número de caracteres, separados por classe (trocadilho e não trocadilho). Observa-se que a maioria das sentenças está concentrada na faixa entre 40 e 100 caracteres, com um pico de frequência próximo aos 50 caracteres para ambas as classes. No gráfico, os textos classificados como trocadilhos são representados em azul, enquanto os não trocadilhos aparecem em verde.

Figura 1 – Histograma de comprimento de texto no Corpus Puntuguese.



A forma geral das distribuições indica uma tendência semelhante entre os dois grupos, embora os textos classificados como trocadilhos apresentem uma ligeira inclinação para comprimentos mais curtos em relação aos não trocadilhos. A frequência diminui gradualmente à medida que o comprimento do texto aumenta, refletindo o padrão esperado para conteúdos breves, como piadas ou perguntas rápidas. Essa análise sugere que o comprimento textual, embora não seja um fator decisivo isoladamente, pode contribuir como feature auxiliar em modelos de classificação automática entre trocadilhos e textos não humorísticos.

O código-fonte utilizado neste trabalho está disponível em: <https://colab.research.google.com/drive/1Vh4vqOC1m1FPCHo3oQ9Lw5jB6hKpTzkK#scrollTo=nZ_VVW4S64en>. O corpus *Puntuguese*, utilizado como base para o treinamento do modelo, está acessível publicamente em: <<https://github.com/Superar/Puntuguese>>.

4 DETALHAMENTO DA ESTRATÉGIA DESENVOLVIDA

Neste artigo, a tarefa principal foi a classificação de trocadilhos utilizando modelos de Processamento de Linguagem Natural (PLN), com ênfase no modelo BERTimbau Large, uma versão treinada do BERT Devlin *et al.* (2019) para a língua portuguesa.

A Figura 2 apresenta uma visão geral do *pipeline* desenvolvido, dividido em sete etapas: coleta do *corpus*, extração de *features*, tokenização, criação do *dataset*, treinamento do modelo, avaliação e visualizações. Essa estrutura modular permitiu o controle e análise sistemática de cada fase, garantindo maior transparência e reprodutibilidade.

Figura 2 – Visão geral das etapas do processo de classificação de trocadilhos utilizando BERTimbau.

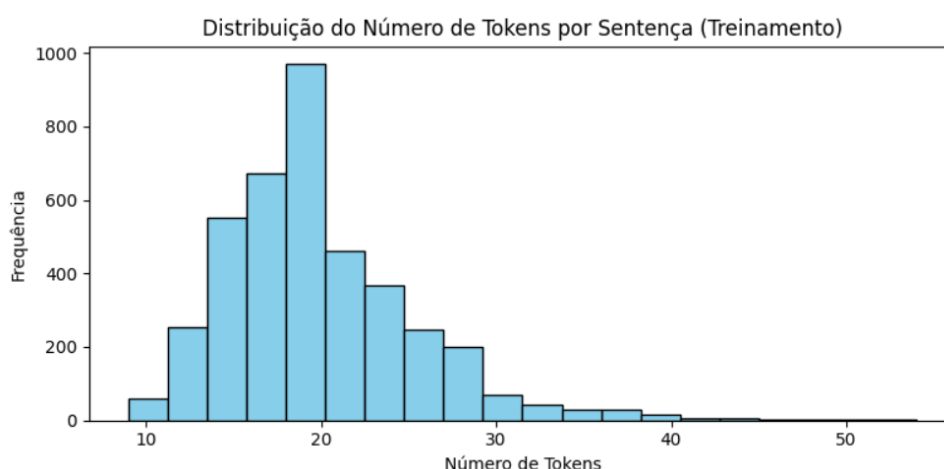


Uma das principais melhorias implementadas foi o ajuste no comprimento máximo das sequências de texto processadas pelo modelo. Inicialmente, o limite era de 50 *tokens*, o que poderia restringir a capacidade do modelo de capturar informações relevantes em sentenças mais longas. Embora as sentenças do corpus possuam, em

média, cerca de 50 caracteres, o número de *tokens* pode ser superior, pois os *tokens* representam não apenas palavras completas, mas também pontuações e fragmentos de palavras.

A Figura 3 apresenta a distribuição do número de *tokens* por sentença no conjunto de treinamento. Observa-se que a maioria das sentenças possui entre 10 e 30 *tokens*, com pico de frequência entre 18 e 20 *tokens*. Esse cenário mostra que a maior parte dos dados está concentrada em sequências curtas, mas também há casos isolados com maior complexidade.

Figura 3 – Distribuição do número de tokens por sentença no conjunto de treinamento.



Apesar de a média estar abaixo de 50 *tokens*, optou-se por aumentar o limite para 250 *tokens*, garantindo que todas as sentenças, inclusive as mais longas e contextualmente ricas, fossem processadas sem truncamento. Essa escolha foi motivada por testes empíricos que demonstraram que o modelo se beneficia de um limite mais amplo, especialmente em trocadilhos que envolvem construções ambíguas, contextos narrativos ou repetições que ampliam o número de tokens mesmo em textos curtos.

Esse ajuste contribuiu para preservar a integridade semântica das sentenças durante o treinamento, o que se refletiu em um ganho de desempenho, permitindo ao modelo capturar melhor as nuances linguísticas e contextuais características dos trocadilhos.

Quanto ao balanceamento das classes, o corpus Puntuguese já era previamente balanceado, com o mesmo número de exemplos para trocadilhos e não trocadilhos. No entanto, para garantir que o modelo tratasse as classes de maneira equitativa durante o treinamento, foi utilizado o cálculo automático dos pesos de classe com base na frequência observada no conjunto de dados. Esses pesos foram aplicados diretamente na função de perda *CrossEntropyLoss* (PyTorch Documentation, 2024), tornando o modelo mais sensível a possíveis erros de classificação em ambas as categorias. A divisão dos dados foi feita da seguinte forma: treinamento: 3.990 amostras (70%), validação: 570 amostras (10%) e teste: 1.140 amostras (20%).

Assim, a proporção adotada foi 70/10/20 para as etapas de treino, validação e teste, respectivamente, assegurando uma separação adequada para ajuste de parâmetros, avaliação intermediária e verificação final de desempenho.

Durante o treinamento, também foram aplicadas estratégias de regularização para evitar *overfitting* e melhorar a capacidade de generalização do modelo. Entre essas estratégias, destaca-se o uso de *dropout*, que ajuda a evitar o sobreajuste ajustando aleatoriamente partes da rede durante o treinamento. Também foi implementado o *early stopping*, que interrompeu o treinamento após três épocas sem melhorias significativas no desempenho, otimizando o uso de recursos computacionais. O treinamento foi realizado ao longo de 6 épocas, proporcionando tempo suficiente para o ajuste eficiente dos parâmetros.

Além disso, foi utilizado um aquecimento (*warmup*) de 100 passos no início do treinamento, o que ajudou a ajustar gradualmente a taxa de aprendizado, definida em $3e-5$. Esse valor possibilitou uma aprendizagem estável, minimizando oscilações bruscas no desempenho do modelo.

As métricas utilizadas para avaliação incluíram *acurácia*, *precisão*, *recall* e *F1-score* ponderado. O uso do *F1-score* ponderado foi especialmente importante para garantir que o modelo mantivesse um bom equilíbrio entre a taxa de acertos e a sensibilidade para cada classe, mesmo diante de pequenas variações em suas distribuições.

Essas melhorias nas estratégias de pré-processamento, ajuste de hiperparâmetros, regularização e avaliação resultaram em uma solução robusta para a tarefa de classificação de trocadilhos. O uso do *BERTimbau Large*, aliado a essas técnicas avançadas, permitiu capturar as sutilezas linguísticas dos trocadilhos, oferecendo uma performance consistente em termos de *acurácia*, *precisão*, *recall* e *F1-score*.

Quadro 1: Configurações do modelo de classificação

Componente	Valor Utilizado
Função de perda	CrossEntropyLoss (com pesos)
Épocas de treino	6
Batch size	16
Learning rate	$3e-5$
Early stopping	Sim (paciência = 3)

Fonte: Elaborado pela autora.

O Quadro 1 apresenta os principais hiperparâmetros empregados no treinamento. A arquitetura adotada incluiu 24 camadas do modelo *BERTimbau Large*, taxa de aprendizado de 3×10^{-5} , *batch size* de 16, função de perda com pesos balanceados (*CrossEntropyLoss*) e estratégia de parada antecipada com paciência igual a 3. Esses parâmetros foram definidos com base em testes preliminares e recomendações da

literatura para modelos baseados em transformadores, como BERT Devlin *et al.* (2019) e BERTimbau Souza, Nogueira & Lotufo (2020).

5 ANÁLISE DOS RESULTADOS ALCANÇADOS

O objetivo principal deste trabalho foi desenvolver um sistema de classificação capaz de identificar automaticamente trocadilhos em textos em português, utilizando o modelo *BERTimbau Large* aliado a estratégias de pré-processamento e extração de padrões linguísticos.

Ao comparar com outros trabalhos na área, o desempenho alcançado neste trabalho é satisfatório, considerando a complexidade da tarefa de detectar trocadilhos, que envolve mais do que simplesmente entender o texto, mas também a interpretação de jogos de palavras e ambiguidades semânticas. Em comparação com o estudo de Inácio *et al.* (2024), que apresentou um *F1-score* de 68,9% utilizando o *corpus Puntuguese*, os resultados deste trabalho são competitivos, considerando que o mesmo conjunto de dados e uma abordagem semelhante foram empregados. A tarefa de detecção de trocadilhos se mantém desafiadora por exigir compreensão contextual e identificação de nuances linguísticas.

A Tabela 1 apresenta o relatório de classificação por classe, com os principais indicadores de desempenho do modelo no conjunto de teste. Observa-se que a classe “Não Trocadilho” obteve maior *recall* (0,81) em comparação à classe “Trocadilho” (0,60), o que reforça a tendência do modelo em identificar com maior sensibilidade textos que não envolvem trocadilhos. Por outro lado, a classe “Trocadilho” apresentou superioridade na precisão (0,76), indicando que, embora o modelo tenha menor sensibilidade para essa classe, tende a cometer menos erros quando a identifica. Tais observações demonstram que o modelo ainda possui margens para aprimoramento no equilíbrio entre as classes, embora o desempenho geral (acurácia = 0,70) seja considerado satisfatório.

Tabela 1 – Relatório de classificação por classe no conjunto de teste

Classe	Precisão	Recall	F1-score	Amostras
Não trocadilho	0.67	0.81	0.73	57
Trocadilho	0.76	0.60	0.67	57
Acurácia	—	—	0.70	114
Média macro	0.71	0.70	0.70	114
Média ponderada	0.71	0.70	0.70	114

Tabela 2 – *

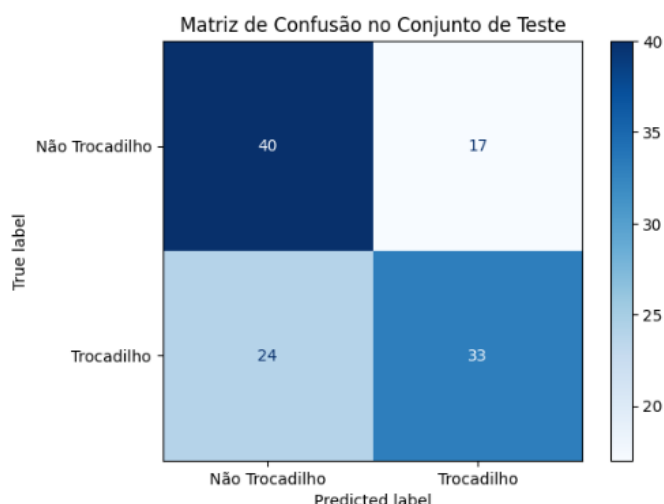
Fonte: Elaborado pela autora.

Além disso, a análise das *features* contextuais utilizadas pelo modelo — especialmente os padrões linguísticos relacionados a trocadilhos, extraídos por expressões

regulares — indica que o modelo foi capaz de capturar nuances importantes do humor verbal. Essas features, como repetições de palavras e terminações específicas (ex.: palavras terminadas em “a” ou “o”), foram extraídas durante o pré-processamento e incorporadas como vetores adicionais no *dataset*. Cada exemplo, ao ser tokenizado, incluía também um vetor com contagens dessas ocorrências, fornecendo informações complementares ao modelo durante o treinamento.

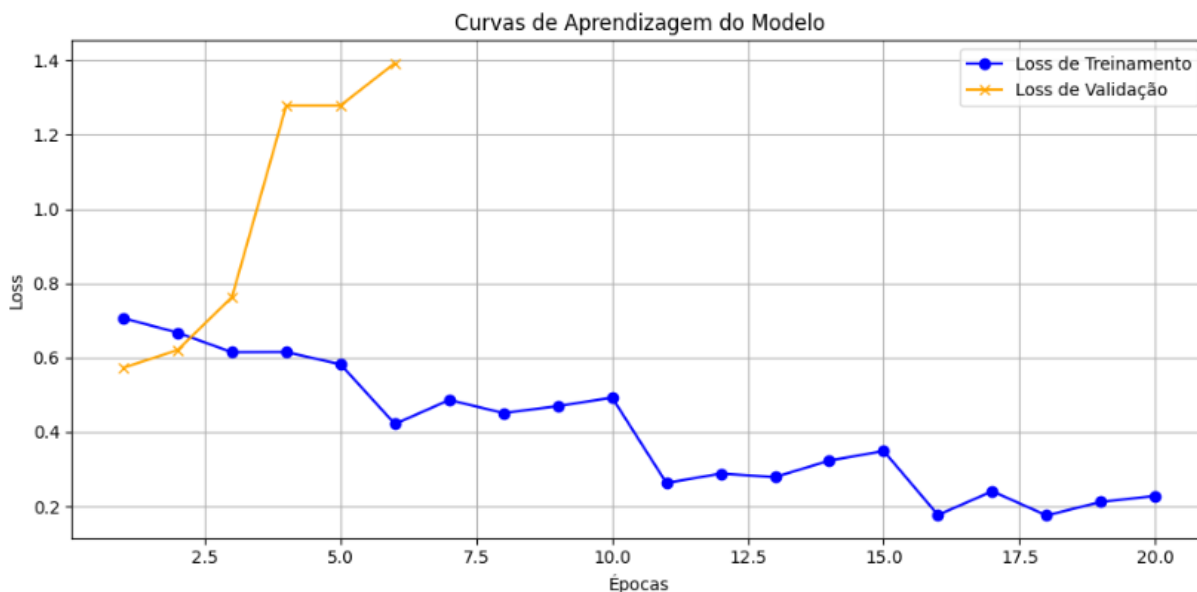
Complementarmente, a Figura 4 apresenta a matriz de confusão do conjunto de teste, evidenciando o desempenho por classe. Observa-se que o modelo apresentou *recall* superior para a classe “Não Trocadilho” (0,81) em comparação à classe “Trocadilho” (0,60), o que indica uma tendência a classificar com maior precisão textos que não envolvem trocadilhos. Ainda assim, o equilíbrio geral das métricas aponta para um desempenho satisfatório.

Figura 4 – Matriz de Confusão no Conjunto de Teste.



Além da avaliação pontual no conjunto de teste, foi realizada uma análise da curva de aprendizagem do modelo, conforme mostrado na Figura 5. A curva revela que a perda de treinamento manteve tendência de queda ao longo das épocas, enquanto a perda de validação começou a subir a partir da quarta época — indício claro de *overfitting*. Essa divergência entre as curvas justifica a adoção da técnica de *early stopping*, que interrompeu o treinamento ao detectar ausência de melhoria na validação por três épocas consecutivas.

Figura 5 – Curvas de Aprendizagem do Modelo (Loss de Treinamento e Validação).



Fonte: Elaborado pela autora.

Essa análise reforça a importância das estratégias de regularização empregadas e da escolha criteriosa do número de épocas de treinamento. O comportamento do modelo, tanto quantitativa quanto qualitativamente, demonstra sua capacidade de identificar trocadilhos de forma robusta, embora com espaço para melhorias em termos de generalização.

Considerações Complementares

Além das representações contextualizadas geradas pelo *BERTimbau Large*, o sistema desenvolvido também incorporou um conjunto de *features* linguísticas extraídas por expressões regulares, como repetições de palavras, uso de pronomes interrogativos, terminações fonológicas específicas e padrões de estrutura frasal. Embora não tenha sido realizado um experimento de ablação para isolar o impacto dessas informações, observou-se durante o desenvolvimento que sua inclusão aumentou a estabilidade das previsões, especialmente em textos curtos e ambíguos. Essas *features* funcionam como complemento ao *embedding* textual, fornecendo pistas adicionais ao classificador.

Adicionalmente, foi conduzida uma análise qualitativa dos erros cometidos pelo modelo no conjunto de teste. A seguir, são apresentados dois exemplos de classificações incorretas:

- **Texto:** “Por que o casal de patos sempre tomam banho quando discutem? Para deixar tudo em pratos limpos.”
Classe verdadeira: Não trocadilho **Classe prevista:** Trocadilho

- **Texto:** “Qual é a atriz mais dramática do mundo? A tristeza.”

Classe verdadeira: Trocadilho

Classe prevista: Não trocadilho

Esses casos ilustram a dificuldade do modelo em lidar com jogos de palavras mais sutis, bem como sua tendência a superestimar o humor em perguntas estruturadas. O uso de construções ambíguas parece induzir o modelo a inferir incorretamente a presença de trocadilhos, mesmo quando a resposta é literal ou não contém ambiguidade semântica. Isso reforça a importância de estratégias adicionais que permitam distinguir trocadilhos implícitos de frases que apenas se assemelham ao estilo de piadas.

6 CONCLUSÃO

Neste artigo, o modelo BERTimbau Large foi aplicado à tarefa de classificação de sentenças com o objetivo de identificar trocadilhos em português. O modelo demonstrou desempenho competitivo, com bons resultados nas primeiras épocas de treinamento, reforçando seu potencial para lidar com esse tipo específico de humor verbal.

Apesar disso, observou-se a ocorrência de *overfitting* a partir da segunda época, caracterizada pela redução contínua da perda de treinamento e aumento progressivo da perda de validação. Esse comportamento sugere que o modelo pode ter aprendido padrões específicos dos dados de treino, sem captar plenamente a lógica subjacente aos trocadilhos. Tal limitação reforça a necessidade de estratégias complementares que favoreçam a generalização.

De modo geral, o estudo demonstrou a viabilidade de aplicar modelos de linguagem avançados à tarefa de classificação de trocadilhos, oferecendo uma base sólida para pesquisas futuras. Os resultados obtidos validam a abordagem proposta e revelam oportunidades relevantes de melhoria, especialmente no equilíbrio entre aprendizado e capacidade de generalização.

6.0.1. TRABALHOS FUTUROS

Entre as possibilidades de aprimoramento, destacam-se a incorporação de mais *features* linguísticas contextuais, como estruturas frasais e ambiguidade semântica, além da adoção de técnicas de regularização mais robustas e maior diversidade no corpus de treinamento. A utilização de *embeddings* semânticos mais profundos também pode contribuir para a detecção de nuances presentes nesse tipo de humor.

REFERÊNCIAS

- BONET, H. A.; RINCÓN, A. M.; LÓPEZ, A. M. Detection, classification and quantification of hurtful humor (huhu) on twitter using classical models, ensemble models, and transformers. In: *IberLEF@ SEPLN*. [S.l.: s.n.], 2023. Citado na página 6.
- CRUZ, J.; ELVIRA, L.; TABERNERO, M.; SEGURA-BEDMAR, I. In unity, there is strength: On weighted voting ensembles for hurtful humour detection. In: *IberLEF@ SEPLN*. [S.l.: s.n.], 2023. Citado na página 7.
- DELABASTITA, D. *Wordplay and Translation: Special Issue of 'The Translator' 2/2 1996*. [S.l.]: Routledge, 2016. Citado na página 5.
- DEVLIN, J.; CHANG, M.-W.; LEE, K.; TOUTANOVA, K. Bert: Pre-training of deep bidirectional transformers for language understanding. In: *Proceedings of the 2019 conference of the North American chapter of the association for computational linguistics: human language technologies, volume 1 (long and short papers)*. [S.l.: s.n.], 2019. p. 4171–4186. Citado 2 vezes nas páginas 9 e 12.
- INACIO, M. L.; WICK-PEDRO, G.; RAMISCH, R.; SANTO, L. E.; CHACON, X. S. Q.; SANTOS, R.; SOUSA, R.; ANCHIÊTA, R.; OLIVEIRA, H. G. Puntuguese: A corpus of puns in Portuguese with micro-edits. In: CALZOLARI, N.; KAN, M.-Y.; HOSTE, V.; LENCI, A.; SAKTI, S.; XUE, N. (Ed.). *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*. Torino, Italia: ELRA and ICCL, 2024. p. 13332–13343. Disponível em: <<https://aclanthology.org/2024.lrec-main.1167/>>. Citado 2 vezes nas páginas 5 e 7.
- MILLER, T.; HEMPELMANN, C.; GUREVYCH, I. SemEval-2017 task 7: Detection and interpretation of English puns. In: BETHARD, S.; CARPUAT, M.; APIDIANAKI, M.; MOHAMMAD, S. M.; CER, D.; JURGENS, D. (Ed.). *Proceedings of the 11th International Workshop on Semantic Evaluation (SemEval-2017)*. Vancouver, Canada: Association for Computational Linguistics, 2017. p. 58–68. Disponível em: <<https://aclanthology.org/S17-2005/>>. Citado na página 6.
- PyTorch Documentation. *torch.nn.CrossEntropyLoss*. 2024. <<https://pytorch.org/docs/stable/generated/torch.nn.CrossEntropyLoss.html>>. Acesso em: 6 abr. 2025. Citado na página 10.
- SOUZA, F.; NOGUEIRA, R.; LOTUFO, R. Bertimbau: pretrained bert models for brazilian portuguese. In: SPRINGER. *Brazilian conference on intelligent systems*. [S.l.], 2020. p. 403–417. Citado 2 vezes nas páginas 5 e 12.

AGRADECIMENTOS

A Deus, que foi minha luz nas horas sombrias e meu abrigo nos dias de tempestade, guiando-me com fé e esperança a cada passo desta jornada.

Ao meu orientador Aislan Rafael e ao meu coorientador Rafael Torres, por oferecerem não apenas conhecimento, mas também paciência, confiança e incentivo, conduzindo-me com firmeza e cuidado até a conclusão deste trabalho.

À minha mãe, Elyne, por ser meu exemplo supremo de força e por seu amor incondicional, a verdadeira base de tudo o que sou. Em você sempre encontrei meu porto seguro, o refúgio seguro em meio às tempestades e a alegria genuína nas celebrações. Esta conquista nada mais é que um espelho da mulher incrível que me ensinou a jamais desistir.

Ao meu namorado e grande companheiro, Humberto. Agradeço por sua presença que superou qualquer distância, por sua compreensão que acalmou minhas dúvidas e por ser a voz da confiança que me reerguia sempre que eu pensava em fraquejar. Seu apoio foi fundamental para que eu chegasse até aqui. Por tudo isso, esta conquista carrega a marca indelével do seu amor.

Aos meus amigos e pilares, João Fernandes, Letícia, Kedna, Emileny e Alcilene, meu porto seguro nos dias de tempestade. Agradeço por compartilharem cada ansiedade, por celebrarem cada vitória como se fosse sua e pela confiança que depositaram em mim, mesmo quando eu duvidava. A jornada foi mais leve porque a amizade de vocês foi o alicerce que me permitiu construir este sonho.

À minha avó, Maria de Fátima, que, mesmo sem muitas palavras, mostrou sua alegria com a minha aprovação na faculdade e transformava sua torcida em gestos, cuidando para que eu estivesse sempre a postos para construir meu futuro.

Ao meu avô, José Demontier(em memória), meu eterno porto seguro. Dele, guardo com orgulho o privilégio de ter sido sempre a sua “mimada” e o tesouro do apelido “Jurema”. A melodia desse nome carrega minhas memórias mais doces e, para mim, sempre significará lar, cumplicidade e afeto. É a certeza de um amor que me fortalece e me acompanhará para sempre.

De maneira especial, dedico esta conquista às minhas amadas irmãs. Que a minha jornada até aqui sirva de farol e inspiração, mostrando que nossos sonhos são alcançáveis. Que esta porta que hoje abro possa facilitar a passagem de vocês para um futuro brilhante, pois meu maior desejo é que voem ainda mais alto. Ser um exemplo para vocês é, talvez, o maior título que recebo hoje.