

Inteligência Artificial e Linguagem Natural em Ambientes Digitais: Uma Síntese Sistemática da Literatura

Fabiane Sorbar
Faculdade Donaduzzi
Toledo, Brasil
fabiane.sorbar@bpkedu.com.br

Anabelle Tait
Faculdade Donaduzzi
Toledo, Brasil
anabelle.tait@alunos.bpkedu.com.br

Jorge Luiz Fanas Guedes
Faculdade Donaduzzi
Toledo, Brasil
jorge.guedes@alunos.bpkedu.com.br

Alan Vitor Bento
Faculdade Donaduzzi
Toledo, Brasil
alan.bento@alunos.bpkedu.com.br

Mayumi Hirata Bogoni
Faculdade Donaduzzi
Toledo, Brasil
mayumi.bogoni@alunos.bpkedu.com.br

Murilo Henrique Zils
Faculdade Donaduzzi
Toledo, Brasil
murilo.zils@alunos.bpkedu.com.br

Pedro Henrique dos Santos Leuchs
Faculdade Donaduzzi
Toledo, Brasil
pedro.leuchs@alunos.bpkedu.com.br

Abstract—This paper presents a systematic literature review on the application of Artificial Intelligence (AI) and Machine Learning (ML) techniques in digital environments, focusing on social networks. Supervised and unsupervised algorithms are explored, such as Naive Bayes, K-Nearest Neighbors (KNN), K-Means, BERT with multicriteria, and PINNs, along with the use of entropy to reduce informational uncertainty. The reviewed studies show that combining these methods enhances accuracy, explainability, and scalability, contributing to advances in natural language analysis, digital behavior, and information security.

Keywords—Artificial Intelligence; Machine Learning; NLP; Social Networks; Entropy.

Resumo—Este artigo apresenta uma síntese sistemática da literatura sobre a aplicação de técnicas de Inteligência Artificial (IA) e Aprendizado de Máquina (AM) em ambientes digitais, com foco em redes sociais. São explorados algoritmos supervisionados e não supervisionados, como Naive Bayes, K-Nearest Neighbors (KNN), K-Means, BERT com multicritérios e PINNs, além do uso da entropia na redução da incerteza informacional. Os estudos analisados demonstram que a escolha e combinação adequada dos modelos permitem interpretações mais precisas, explicáveis e escaláveis, promovendo avanços significativos na análise de linguagem natural, comportamento digital e segurança da informação.

Palavras-chave—Inteligência Artificial; Aprendizado de Máquina; PLN; Redes Sociais; Entropia.

I. INTRODUÇÃO

O crescimento acelerado das redes sociais tem gerado volumes massivos de dados que, quando bem explorados, oferecem *insights* valiosos sobre comportamentos, tendências e padrões. Nesse cenário, técnicas de inteligência artificial e aprendizado de máquina são fundamentais para transformar dados brutos em informação útil. O objetivo deste trabalho foi realizar uma comparação entre os modelos utilizados durante o Processamento de Linguagem Natural (PLN) em redes sociais, analisando suas abordagens, aplicações e resultados obtidos.

Modelos como o BERT, aliados a critérios multicritérios, se destacam na interpretação de textos, enquanto modelos mais tradicionais, como Naive Bayes e K-Nearest Neighbors (KNN), seguem sendo eficazes em tarefas de classificação e recomendação. Algoritmos não supervisionados, como o K-Means, auxiliam na segmentação de dados, e conceitos como entropia ajudam na avaliação de incertezas e relevância da informação.

Mais recentemente, abordagens como as Physics-Informed Neural Networks (PINNs) começam a ser aplicadas além do contexto físico, oferecendo novas formas de modelar comportamentos complexos, como a disseminação de informações nas redes sociais. A combinação desses modelos permite uma análise mais precisa e eficiente dos dados,

contribuindo para soluções cada vez mais inteligentes no ecossistema digital.

Para assegurar rigor e transparência na seleção e análise da literatura, a presente pesquisa utilizou parcialmente a metodologia PRISMA (Preferred Reporting Items for Systematic Reviews and Meta-Analyses) para a condução da síntese sistemática, realizando o estudo por meio de buscas em plataformas de pesquisa, com o objetivo de identificar estudos relevantes sobre o tema.

II. INTELIGÊNCIA ARTIFICIAL E O APRENDIZADO DE MÁQUINA

A inteligência artificial (IA) é um ramo da ciência da computação voltado ao desenvolvimento de sistemas que simulam capacidades cognitivas humanas, como raciocínio, aprendizado e tomada de decisão [1]. Um de seus principais subcampos, o Aprendizado de Máquina (AM), permite que sistemas ajustem seus modelos a partir de dados, podendo ser supervisionado, quando utiliza dados rotulados, ou não supervisionado, que busca padrões sem rótulos prévios [2]. Essa abordagem possibilita tarefas complexas, como reconhecimento de linguagem natural e previsão de comportamentos coletivos.

Nas redes sociais, a aplicação de IA enfrenta desafios devido à heterogeneidade dos dados — que incluem texto, imagens, vídeos e interações — e à necessidade de processamento em tempo real frente à velocidade e informalidade do conteúdo [3]. Algoritmos supervisionados como Naive Bayes e K-Nearest Neighbors (KNN) são comuns em classificação de sentimentos e análise de usuários [4], enquanto modelos não supervisionados como K-means identificam padrões latentes em grandes volumes de dados textuais [5]. Abordagens avançadas, como Physics-Informed Neural Networks (PINNs) e modelos baseados em BERT, combinam conhecimento de domínio com capacidades avançadas de processamento de linguagem natural [6], [7].

Este artigo explora o uso desses algoritmos e técnicas de Processamento de Linguagem Natural (PLN) em estudos empíricos, destacando seu potencial para compreender a dinâmica das redes sociais e seu impacto em comunicação, marketing, psicologia e análise de tendências sociais.

III. ALGORITMOS IMPORTANTES PARA O APRENDIZADO DE MÁQUINA

A análise de dados de redes sociais para compreender comportamentos, tendências e interações online exige uma gama diversificada de algoritmos de Aprendizado de Máquina (AM) e técnicas de Processamento de Linguagem Natural (PLN). A escolha do algoritmo adequado depende da natureza dos dados e dos objetivos do estudo, variando desde a

classificação de sentimentos e a detecção de desinformação até a predição de padrões complexos.

Iremos explorar alguns dos algoritmos e abordagens mais relevantes que têm sido aplicados em experimentos e estudos de caso para extrair insights valiosos das plataformas digitais.

A. Physics-informed neural networks (PINNs)

As Physics-Informed Neural Networks (PINNs) integram o conhecimento prévio das leis da física às redes neurais artificiais. Essa técnica tem se destacado por sua capacidade de modelar fenômenos complexos em diferentes domínios, especialmente onde dados empíricos são limitados, mas há conhecimento consolidado sobre as equações governantes do sistema.

Nos últimos anos, estudos têm explorado o potencial das PINNs em contextos sociais e digitais, como demonstrado em diversos trabalhos analisados nesta síntese. Por exemplo, um estudo de 2024 propôs um modelo baseado em PINNs para prever o crescimento de usuários em plataformas de redes sociais, demonstrando resultados promissores na predição de padrões de crescimento a partir de uma abordagem fundamentada em leis físicas do sistema dinâmico [8].

Outros trabalhos ampliaram o uso das PINNs para além da dinâmica populacional, como no estudo de opinião pública. Em 2022, pesquisadores estenderam esse conceito ao desenvolver redes neurais informadas sociologicamente (SINNs), que integram aspectos das PINNs com dinâmicas sociais para prever a evolução das opiniões de usuários em redes como Twitter e Reddit. O modelo superou modelos tradicionais, indicando que a incorporação de estruturas sociais e físicas pode enriquecer previsões em ambientes digitais altamente complexos [9].

No contexto da saúde pública, um estudo de 2023 empregou PINNs para investigar a propagação da COVID-19 na China após o relaxamento das medidas sanitárias. O modelo foi capaz de estimar taxas reais de infecção e quantificar os efeitos da subnotificação, da eficiência vacinal e do comportamento social, evidenciando a robustez das PINNs em lidar com dados incompletos e heterogêneos em problemas epidemiológicos [10].

Mais recentemente, em 2025, uma pesquisa combinou PINNs com mecânica lagrangiana para identificar dinâmicas migratórias entre regiões da Rússia e dos Estados Unidos. O modelo conseguiu prever fluxos migratórios e avaliar os efeitos de políticas públicas sobre padrões de deslocamento populacional, superando abordagens tradicionais baseadas exclusivamente em indicadores macroeconômicos. Esse trabalho ressalta o potencial das PINNs para aplicações em ciências sociais, planejamento urbano e políticas públicas inteligentes [11].

Esses exemplos demonstram que as PINNs não apenas ampliam as fronteiras do aprendizado de máquina, mas também fortalecem a capacidade de modelagem em cenários onde o conhecimento físico e social é essencial. Assim, as PINNs emergem como uma ferramenta versátil e poderosa para a modelagem de sistemas complexos em áreas que vão muito além das ciências exatas, alcançando com eficácia domínios sociais, epidemiológicos e digitais.

B. Naive Bayes

Algoritmos supervisionados de aprendizado de máquina, como o Naive Bayes, têm sido amplamente utilizados na análise de sentimentos em contextos sociais e políticos, especialmente com grandes volumes de dados textuais provenientes de redes sociais. Baseado no teorema de Bayes e assumindo independência condicional entre as características, o Naive Bayes é reconhecido por sua simplicidade, eficiência computacional e boa performance em tarefas de classificação textual, mesmo diante da linguagem informal e ruidosa típica desses ambientes digitais. Essa abordagem se mostra capaz de lidar com a complexidade dos dados, oferecendo resultados rápidos e interpretáveis que contribuem para a compreensão das opiniões e sentimentos expressos online [12].

Em uma das pesquisas estudadas, o Naive Bayes foi utilizado para observar como consumidores se posicionam diante de marcas que assumem compromissos com causas sociais — como sustentabilidade, empoderamento feminino ou temas ligados à cultura brasileira. O estudo analisou tweets de usuários, buscando identificar automaticamente sentimentos e emoções relacionados às marcas. Segundo Qi e Shabrina [12], o algoritmo mostrou-se eficiente nessa tarefa, muito por conta da sua simplicidade e desempenho aceitável em classificações de texto. Apesar de não ser o mais sofisticado, ele se saiu bem, principalmente pela capacidade de lidar com linguagem informal. Os autores destacam que a abordagem de aprendizado supervisionado, como o Naive Bayes, é capaz de lidar com grandes volumes de dados textuais, mesmo com linguagem espontânea, apresentando resultados satisfatórios na análise de sentimentos [12]. Os resultados deram indícios de que modelos como esse podem ajudar a entender melhor a percepção do público e, com isso, embasar estratégias de comunicação mais sensíveis ao sentimento coletivo.

O segundo estudo partiu para um campo diferente: a política. Nele, o Naive Bayes foi aplicado à análise de cerca de 3.000 tweets sobre as eleições presidenciais da Indonésia em 2024. A ideia era mapear o sentimento do público em relação aos principais candidatos — se positivo, negativo ou neutro. De acordo com Hakiki, Pambudi e Asriyanik [13], o modelo obteve uma acurácia média de até 88% na classificação dos sentimentos, dependendo do candidato analisado. O estudo

reforça o potencial da mineração de dados em tempo real para a leitura de tendências políticas, principalmente em redes como o X (antigo Twitter), onde o volume e a espontaneidade das opiniões são significativos [13].

Juntos, os dois estudos demonstram que, mesmo sendo um modelo relativamente simples, o Naive Bayes pode dar conta de análises de sentimento em ambientes com muitos dados e linguagem informal. Sua transparência, leveza e baixo custo computacional o tornam adequado para pesquisas sociais, onde nem sempre é necessário o uso de modelos muito pesados ou difíceis de interpretar [12], [13].

Ao trazer exemplos que vão do consumo à política, este trabalho mostra como técnicas de Processamento de Linguagem Natural combinadas com modelos estatísticos podem ajudar a entender fenômenos sociais mais amplos. É uma forma de tornar esse tipo de análise mais acessível, rápida e útil para diferentes áreas.

C. K-Nearest Neighbors - KNN

O K-Nearest Neighbors (KNN) é um algoritmo de aprendizado supervisionado conhecido por sua simplicidade e interpretabilidade inerente. Baseado no princípio da “votação dos vizinhos mais próximos”, o KNN classifica uma amostra com base nas categorias dos dados de treinamento mais próximos no espaço de características. Essa transparência na tomada de decisão permite rastrear diretamente os resultados até exemplos específicos, característica fundamental para a inteligência artificial explicável (XAI). Por essa razão, o KNN é valorizado em áreas que demandam confiança, auditabilidade e clareza, especialmente quando comparado a modelos mais complexos e opacos, frequentemente chamados de “caixas-pretas” [14].

Pensando nisso, no estudo desenvolvido por Hazard et al. [15], os autores exploram a ideia de que a interpretabilidade, ou seja, a capacidade de entender como e por que um modelo de IA toma determinada decisão, deve ser parte central da construção dos algoritmos, e não algo adicionado depois. Então, utilizando o KNN como exemplo, eles mostram como certos algoritmos já são naturalmente mais fáceis de interpretar, pois suas decisões são baseadas diretamente nos dados mais próximos no espaço de treinamento. Eles propõem ir além: integrar essa lógica intuitiva com ferramentas da Teoria da Informação, que ajudam a medir e explicar a quantidade de “certeza” ou “incerteza” em uma decisão. Essa proposta é aplicada em contextos chamados de modelos targetless, onde o sistema aprende sem uma resposta certa definida previamente — o que exige ainda mais clareza sobre como o modelo raciocina. O objetivo final é desenvolver sistemas de IA que sejam não só eficazes, mas também confiáveis, transparentes e compreensíveis desde o início.

Contudo, sabemos que um dos principais desafios das técnicas de machine learning quando aplicadas a grandes volumes de dados é o seu alto custo computacional, já que precisam comparar cada nova entrada com todos os dados de treinamento, o que pode ser muito lento. Para resolver isso, um dos estudos analisados apresenta uma solução que utiliza GPUs, que são processadores capazes de realizar muitas operações ao mesmo tempo. A novidade do trabalho está em agrupar várias consultas diferentes e enviá-las juntas para a GPU, permitindo que o processamento aconteça de forma paralela e mais eficiente. Isso reduz o tempo gasto com a transferência dos dados e o overhead de cada consulta individual. Como resultado, o algoritmo consegue realizar bilhões de comparações por segundo, tornando o KNN muito mais rápido e escalável para aplicações que precisam de respostas em tempo real e que trabalham com grandes bases de dados [16].

Além disso, também existem estudos promissores em que foi combinado o KNN com técnicas modernas de aprendizado profundo para potencializar seus pontos fortes. No estudo de Arora et al. [17] eles desenvolveram um modelo que une a simplicidade e transparência do KNN com a capacidade das redes neurais de extrair características complexas de imagens. As redes neurais são ótimas em aprender representações complexas e sutis de dados, especialmente em imagens, mas às vezes suas decisões podem parecer difíceis de entender. Já o KNN, por sua vez, é simples e transparente, pois baseia suas decisões em exemplos próximos e conhecidos. Ao juntar esses dois métodos, o modelo consegue extrair características detalhadas das imagens com a rede neural, e depois usa o KNN para tomar decisões mais claras e explicáveis. Isso melhora não só a precisão da classificação, mas também a confiança no resultado, algo fundamental em aplicações reais que precisam de rapidez e segurança, como sistemas de reconhecimento facial e diagnóstico por imagem.

Esses estudos analisados deixam claro que o K-Nearest Neighbors - KNN possui uma capacidade natural de explicar decisões com base em exemplos reais tornando-o uma escolha valiosa em áreas onde a confiança e a compreensão são essenciais. Além disso, ao integrar avanços como o uso de GPUs para acelerar o processamento e a combinação com redes neurais para extrair características mais complexas, o KNN se mostra pronto para enfrentar os desafios dos grandes volumes de dados e das aplicações modernas. Assim, o KNN continua sendo uma ferramenta versátil, capaz de unir a simplicidade da lógica intuitiva com a potência da tecnologia atual.

D. Bert (Multicritérios)

Abordagens baseadas em análise multicritério (MCDA) e processamento de linguagem natural (PLN) têm se mostrado

eficazes na resolução de problemas complexos e com múltiplas variáveis em contextos diversos, como a indústria, a educação infantil e os sistemas jurídicos. Nesse cenário, destaca-se o uso do BERT (Bidirectional Encoder Representations from Transformers), modelo pré-treinado desenvolvido pela Google que revolucionou o PLN ao permitir a interpretação bidirecional do contexto das palavras em uma sentença. Sua capacidade de capturar relações semânticas profundas torna-o especialmente eficaz para tarefas como classificação de texto, extração de informação e análise de sentimentos. Ao ser integrado a sistemas de apoio à decisão multicritério, o BERT potencializa a compreensão de textos complexos, contribuindo para decisões mais precisas e contextualizadas em ambientes com alta demanda informacional.

No campo da otimização industrial, o uso dos métodos AHP e TOPSIS no problema de corte de estoque mostrou que decisões mais equilibradas e eficientes podem ser alcançadas ao considerar múltiplos critérios simultaneamente. A solução proposta teve excelente desempenho em termos de aproveitamento de material, com baixo índice de perda (GAP inferior a 5%), provando que a análise multicritério pode complementar (e até superar) modelos tradicionais de programação inteira ao oferecer mais flexibilidade no processo decisório [18].

No contexto da educação infantil, a aplicação da teoria MAUT na criação da plataforma MAUT-I e de um aplicativo para avaliação de mídias interativas trouxe uma inovação importante: a padronização do julgamento qualitativo sobre conteúdos voltados à primeira infância. Ao permitir que critérios como interatividade, segurança e adequação pedagógica fossem ponderados de maneira estruturada, foi possível gerar um índice único de qualidade, que auxilia pais, educadores e desenvolvedores na escolha de ferramentas digitais mais apropriadas para crianças pequenas [19].

Avançando para o domínio jurídico, os dois outros estudos analisados mostram o impacto do PLN aliado a modelos de aprendizado profundo. O primeiro utilizou o modelo LEGAL-BERT, ajustado ao domínio legal brasileiro, para realizar modelagem de tópicos em decisões judiciais, alcançando alto grau de coerência temática com avaliações humanas (acurácia temática de 84,6%), o que demonstra a viabilidade do uso de modelos linguísticos especializados para apoiar a sistematização e a análise documental em larga escala [20]. O segundo artigo propôs um método alternativo baseado em subespaços vetoriais (ACDSLQV) para classificação de decisões judiciais relacionadas à habitação, atingindo acurácia de 93% com custo computacional reduzido e alta interpretabilidade — um diferencial essencial em ambientes regulatórios sensíveis [21].

De forma integrada, esses quatro trabalhos demonstram como

técnicas quantitativas e computacionais podem ser adaptadas para auxiliar a tomada de decisão em contextos onde múltiplos fatores precisam ser considerados simultaneamente, sejam eles operacionais, pedagógicos ou legais. Além disso, reforçam a importância de desenvolver ferramentas acessíveis — como plataformas, aplicativos e modelos interpretáveis — que não apenas entreguem resultados eficientes, mas também ampliem a compreensão dos usuários sobre os critérios envolvidos nas decisões.

Ao reunir MCDA e PLN em diferentes campos, esta pesquisa abre caminho para futuras aplicações interdisciplinares, onde a combinação entre inteligência artificial e métodos estruturados de decisão pode promover soluções mais justas, transparentes e eficazes.

E. Entropia

Após analisar alguns estudos recentes focados em mídias sociais — um voltado para a predição de personalidade, outro para a detecção de desinformação — fica claro que ambos se apoiam em técnicas avançadas de machine learning e processamento multimodal de dados. Isso reforça o quanto a inteligência artificial tem potencial para entender e analisar comportamentos nessas plataformas. Um ponto que chama atenção é o papel da entropia, conceito essencial na teoria da informação e na IA, que mede a incerteza ou aleatoriedade dos dados. Diminuir essa entropia é uma tarefa central para muitos algoritmos de aprendizado, pois assim eles conseguem lidar melhor com a complexidade e a imprevisibilidade que caracterizam as informações das redes sociais, transformando dados crus em previsões e classificações confiáveis [22].

No primeiro estudo, foi criado um modelo para prever traços de personalidade a partir dos dados coletados em redes sociais. A metodologia combinou diversas técnicas de machine learning, como o Binary-Partitioning Transformer (BPT) junto com Term Frequency & Inverse Gravity Moment (TF-IGM), e diversas arquiteturas de deep learning — incluindo redes neurais recorrentes (RNNs), Long Short-Term Memory (LSTM), redes convolucionais (CNNs), K-Means, GRU, autoencoders, além de modelos como BERT, RoBERTa e XLNet. Ao processar dados do Facebook, Twitter e Instagram, o modelo conseguiu reduzir a entropia dos atributos dos usuários para classificá-los em perfis de personalidade, alcançando resultados expressivos: F1-score de 0,762 e precisão de 78,34% no Facebook; 0,783 e 79,67% no Twitter; e 0,821 e 86,84% no Instagram [23].

Outro trabalho, de Joshi et al. [24], propõe um modelo inovador para identificar desinformação, com um olhar especial para a explicabilidade das decisões. A arquitetura multimodal combina dados de texto e imagem: o módulo de texto usa Transformers pré-treinados, como BERT, enquanto o módulo de

imagem emprega redes convolucionais. Antes da classificação, esses dados passam por um módulo de fusão multimodal. O diferencial está no uso de técnicas explicativas (XAI), como LIME e SHAP, que apontam quais palavras ou regiões das imagens influenciaram a decisão do sistema. Isso permite diminuir a entropia do conteúdo em várias plataformas — Twitter, Facebook e Instagram — classificando com alta confiança o que é informação verdadeira ou falsa, e ainda explicando o motivo da decisão.

Além disso, a relevância da entropia para machine learning é bem destacada por outras fontes que tratam do tema, mostrando como esse conceito está presente em diversas técnicas, por exemplo, na construção de árvores de decisão ou nas funções de perda de redes neurais. Entender e controlar a entropia é, portanto, vital para lidar com dados complexos e incertos, como os extraídos das mídias sociais [25].

De forma geral, esses estudos se complementam ao mostrar como técnicas de aprendizado avançadas buscam reduzir a incerteza nos dados para extrair padrões significativos. A utilização do processamento multimodal e a análise em várias plataformas reforçam as possibilidades práticas desses trabalhos, que vão desde personalização em marketing e análise detalhada de consumidores até a detecção de fake news, desenvolvimento de interfaces mais adaptativas, pesquisas psicológicas e maior transparência em sistemas de IA.

Mas é claro que ainda existem desafios. Entre eles, destacam-se a necessidade de validar os modelos com conjuntos maiores e mais diversos, a dificuldade em anotar dados com qualidade, questões de ética e privacidade, e a constante necessidade de atualizar os algoritmos para acompanhar as mudanças rápidas da informação online. Outro ponto crucial é garantir que as decisões dos sistemas sejam interpretáveis e explicáveis para ganhar a confiança dos usuários.

Por fim, pode-se dizer que a combinação dessas pesquisas representa um avanço importante na análise de mídias sociais. Ao juntar técnicas sofisticadas de IA com preocupações éticas e foco na explicabilidade, abre-se caminho para compreender melhor o comportamento humano e enfrentar a desinformação. Controlar a entropia mostra-se, assim, essencial para superar a incerteza inerente à informação digital.

F. K-Means

O K-means é um algoritmo de aprendizado de máquina não supervisionado amplamente utilizado na tarefa de agrupamento (clusterização) de dados, sendo valorizado por sua simplicidade, eficiência computacional e capacidade de lidar com grandes volumes de informações em alta dimensionalidade. Em cenários em que os dados não possuem rótulos prévios, o K-means torna-se especialmente útil para a descoberta de padrões e temas recorrentes de forma

automatizada. Este estudo aplica o algoritmo na análise de dados textuais extraídos da Dark Web, com o objetivo de identificar agrupamentos naturais e inferir categorias de conteúdo malicioso ou relevante para a segurança digital.

Conforme apontado por Jesus Filho [26] e Lopes [27], o algoritmo foi utilizado para examinar postagens extraídas de fóruns ilegais, como Hidden Answers e Deep Answers, após rigorosa filtragem baseada na relevância dos termos por meio da técnica TF-IDF. A partir da matriz vetorial resultante, foram testadas diferentes quantidades de clusters ($k = 2, 3, 4$ e 5), visando alcançar uma segmentação equilibrada entre clareza temática e diversidade de conteúdo.

A escolha do K-means revelou-se adequada em virtude de sua simplicidade, eficiência computacional e capacidade de lidar com dados de alta dimensionalidade. Quando combinado a técnicas de redução de dimensionalidade, como a Análise de Componentes Principais (PCA), o algoritmo possibilitou a visualização bidimensional dos agrupamentos, o que facilitou a interpretação dos resultados.

Apesar das vantagens, o K-means apresenta limitações importantes. A necessidade de definir previamente o número de clusters (k) pode comprometer a qualidade da segmentação, especialmente se essa escolha for inadequada. Em testes com $k = 5$, por exemplo, observou-se uma especificação excessiva dos dados, dificultando a distinção clara entre alguns grupos, principalmente os concentrados no centro dos gráficos de dispersão.

Outra limitação está na sensibilidade do algoritmo a valores atípicos e à inicialização aleatória dos centróides, o que pode resultar em convergências subótimas. Entretanto, essas fragilidades podem ser minimizadas com ajustes prévios nos dados e a adoção de abordagens complementares, como a modelagem de tópicos via LDA (Latent Dirichlet Allocation).

Em síntese, os estudos de Jesus Filho [26] e Lopes [27] destacam o K-means como uma ferramenta robusta e acessível para análise de agrupamentos em cenários com grande volume de dados textuais e linguagem informal. Apesar de sua simplicidade em relação a algoritmos mais complexos, sua eficiência e facilidade de implementação o tornam adequado a aplicações práticas que demandam extração rápida e escalável de conhecimento.

Ao integrar técnicas de Processamento de Linguagem Natural (PLN) com modelos de clusterização, este trabalho evidencia como a análise automatizada pode apoiar a compreensão de fenômenos digitais complexos — desde a categorização de discussões sociais até a detecção de conteúdos ilícitos — promovendo investigações mais ágeis, precisas e fundamentadas em dados reais.

IV. METODOLOGIA

Este estudo se enquadra de forma parcial na metodologia PRISMA que é amplamente utilizada para garantir rigor e transparência em revisões sistemáticas. Segundo Moher et al. [28], a metodologia PRISMA estabelece diretrizes claras para a identificação, seleção e inclusão de estudos relevantes, promovendo maior qualidade e replicabilidade nas revisões sistemáticas.

As estratégias de busca deste trabalho foram aplicadas em plataformas reconhecidas, como Google Acadêmico, Portal de Periódicos da CAPES e Scopus, com critérios de inclusão e exclusão previamente definidos, priorizando artigos recentes e relevantes ao tema de Processamento de Linguagem Natural (PLN) em redes sociais. A triagem inicial considerou o ano de publicação e a pertinência dos trabalhos ao objetivo da pesquisa.

Cabe salientar que a metodologia PRISMA foi empregada de maneira parcial, considerando as especificidades do estudo, visto que o artigo foi desenvolvido por um grupo de acadêmicos, em que cada integrante pesquisou sobre um modelo específico adicionando “redes sociais” e “PLN” como palavras-chave. Dessa forma, não foi possível contabilizar o número de artigos encontrados e excluídos como pede a metodologia.

V. RESULTADOS E DISCUSSÕES

A análise comparativa dos algoritmos evidencia que cada modelo apresenta vantagens específicas, sendo que a escolha do mesmo fica dependente do contexto, do tipo de dados e dos objetivos da aplicação. As PINNs demonstram potencial para modelar dinâmicas sociais complexas ao incorporar restrições físicas e sociais, sendo promissoras em contextos com dados escassos. O Naive Bayes mostrou-se eficaz para análise de sentimentos em ambientes digitais, destacando-se pela leveza computacional e facilidade de interpretação. O KNN mantém relevância pela transparência e interpretabilidade, especialmente quando aliado ao uso de GPUs e extração de atributos por redes neurais. O BERT, associado a abordagens multicritério, provou ser versátil e robusto em tarefas de PLN em diversos setores, aliando profundidade analítica e acessibilidade. A entropia, por sua vez, surge como métrica relevante no controle da incerteza e na detecção de desinformação. Já o K-Means demonstrou utilidade na segmentação de dados não rotulados, apesar de limitações como a sensibilidade a *outliers*. Os resultados reforçam que soluções mais eficazes em redes sociais digitais decorrem da combinação estratégica entre algoritmos, respeitando critérios como volume, explicabilidade e finalidade da análise.

A Tabela I apresenta um quadro comparativo entre os algoritmos / modelos abordados na pesquisa, destacando

TABELA I
QUADRO COMPARATIVO DOS ALGORITMOS CITADOS

Modelo	Pontos Fortes	Limitações
Physics-Informed Neural Networks (PINNs)	Integra conhecimento físico e social. Modela sistemas complexos com dados escassos.	Alta complexidade computacional. Requer conhecimento especializado para implementação.
Naive Bayes	Simplicidade e eficiência computacional. Bom desempenho em classificação textual.	Assunção de independência condicional pode ser irrealista. Menor precisão em dados complexos.
K-Nearest Neighbors (KNN)	Transparência e interpretabilidade. Fácil de implementar.	Alto custo computacional em grandes volumes de dados. Sensível a outliers.
BERT (Multicritérios)	Captura relações semânticas profundas. Versátil e robusto em tarefas de PLN.	Requer alto poder computacional. Complexidade na interpretação dos resultados.
Entropia	Reduz incerteza em dados complexos. Útil para análise de desinformação.	Difícil de aplicar em dados não estruturados. Requer técnicas complementares para maior eficácia.
K-Means	Simplicidade e eficiência em grandes volumes de dados. Útil para segmentação inicial.	Necessidade de definir o número de clusters (k). Sensível a outliers e inicialização aleatória.

seus principais pontos fortes e limitações no contexto do Processamento de Linguagem Natural (PLN) aplicado a redes sociais. Essa comparação permitiu observar que cada modelo possui características específicas quanto à complexidade, interpretabilidade e aplicabilidade, variando conforme o tipo de dado e o objetivo da análise. A partir dessa análise comparativa, foi possível compreender de forma mais ampla as vantagens e desafios envolvidos na escolha do algoritmo mais adequado para diferentes cenários de estudo.

VI. CONCLUSÃO

A análise evidencia o avanço significativo da inteligência artificial nas redes sociais, destacando como diferentes algoritmos e modelos oferecem perspectivas complementares na compreensão do comportamento digital. A orquestração de modelos combinando robustez, inovação e especialização, mostra-se mais valiosa do que a busca por um algoritmo único e superior. Em síntese, o valor reside menos na perfeição algorítmica e mais no uso estratégico e ético das ferramentas disponíveis para compreender e melhorar as interações digitais.

Dado o grande volume de dados gerados continuamente pelas redes sociais, é essencial que os algoritmos utilizados sejam eficientes e rápidos, capazes de lidar com a demanda de processamento em tempo real. Nesse contexto, este estudo

sugere a exploração de combinações estratégicas de algoritmos que possam atender às necessidades de diferentes segmentos, equilibrando precisão, escalabilidade e explicabilidade. A integração de técnicas complementares pode potencializar a análise de dados, permitindo soluções mais robustas e adaptadas às dinâmicas complexas dos ambientes digitais.

Pretende-se, em trabalhos futuros, aprofundar os conhecimentos sobre os modelos estudados, buscando uma combinação de sucesso entre eles, considerando que cada algoritmo possui características próprias, com prós e contras que influenciam diretamente sua aplicação e desempenho.

REFERÊNCIAS

- [1] S. J. Russell and P. Norvig, *Artificial Intelligence: A Modern Approach*, 3rd ed. Prentice Hall, 2010, acesso em 20 junho 2025. [Online]. Available: <https://people.engr.tamu.edu/guni/csce625/slides/AI.pdf>
- [2] C. M. Bishop, *Pattern Recognition and Machine Learning*. Springer, 2006, acesso em 20 junho 2025. [Online]. Available: <https://www.microsoft.com/en-us/research/wp-content/uploads/2006/01/Bishop-Pattern-Recognition-and-Machine-Learning-2006.pdf>
- [3] R. Zafarani, M. A. Abbasi, and H. Liu, *Social Media Mining: An Introduction*. Cambridge University Press, 2014, acesso em 20 junho 2025. [Online]. Available: https://assets.cambridge.org/9781107018853/frontmatter/9781107018853_frontmatter.pdf
- [4] B. Liu, *Sentiment Analysis and Opinion Mining*. Morgan & Claypool Publishers, 2012, acesso em 20 junho 2025. [Online]. Available: <https://www.cs.uic.edu/~liub/FBS/SentimentAnalysis-and-OpinionMining.pdf>
- [5] C. C. Aggarwal and C. Zhai, *Mining Text Data*. Springer, 2012, acesso em 20 junho 2025. [Online]. Available: https://www.researchgate.net/publication/268183036_Mining_Text_Data
- [6] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "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*, vol. 1, 2019, pp. 4171–4186, acesso em 20 junho 2025. [Online]. Available: <https://aclanthology.org/N19-1423.pdf>
- [7] M. Raissi, P. Perdikaris, and G. E. Karniadakis, "Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations," *Journal of Computational Physics*, vol. 378, pp. 686–707, 2019, acesso em 20 junho 2025. [Online]. Available: <https://arxiv.org/pdf/1711.10561>
- [8] L. Kong, R. Z. Shi, and M. Wang, "A physics-informed neural network model for social media user growth," *AIMS Press*, 2024, acesso em 25 maio 2025. [Online]. Available: <https://www.aimspress.com/aimspress-data/aci/2024/2/PDF/aci-04-02-012.pdf>
- [9] M. Okawa and T. Iwata, "Predicting opinion dynamics via sociologically-informed neural networks," *arXiv*, 2022, acesso em 25 maio 2025. [Online]. Available: <https://arxiv.org/pdf/2207.03990>
- [10] S. Ghosh, A. Oueda-Oliva, A. Ghosh, M. Banerjee, and P. Seshaiyer, "Understanding the implications of under-reporting, vaccine efficiency and social behavior on the post-pandemic spread using physics informed neural networks: A case study of china," *PLOS ONE*, vol. 18, no. 8, 2023, acesso em 25 maio 2025. [Online]. Available: <https://journals.plos.org/plosone/article?id=10.1371/journal.pone.0290368>
- [11] K. Zakharov, A. Kovantsev, and A. Boukhanovsky, "Coupling of lagrangian mechanics and physics-informed neural networks for the identification of migration dynamics," *ResearchGate*, 2025, acesso em 25 maio 2025. [Online]. Available: https://www.researchgate.net/publication/389652645_Coupling_of_Lagrangian_Mechanics_and_Physics-Informed_Neural_Networks_for_the_Identification_of_Migration_Dynamics

- [12] Y. Qi and Z. Shabrina, "Sentiment analysis using twitter data: A comparative application of lexicon- and machine-learning-based approach," *Social Network Analysis and Mining*, vol. 13, no. 1, 2023, acesso em 14 maio 2025. [Online]. Available: <https://doi.org/10.1007/s13278-023-01030-x>
- [13] R. Hakiki, A. Pambudi, and S. Asriyanik, "Classification of public sentiment toward 2024 presidential candidates on social media platform x using naïve bayes algorithm," *Journal of Artificial Intelligence and Engineering Applications (JAIEA)*, vol. 3, no. 2, pp. 551–556, fev. 2024, acesso em 14 junho 2025. [Online]. Available: <https://doi.org/10.59934/jaiea.v3i2.422>
- [14] R. K. Halder *et al.*, "Enhancing k-nearest neighbor algorithm: a comprehensive review and performance analysis of modifications," *Journal of Big Data*, vol. 11, p. 113, 2024, acesso em 14 maio 2025. [Online]. Available: <https://doi.org/10.1186/s40537-024-00973-y>
- [15] J. Hazard *et al.*, "Natively interpretable machine learning and artificial intelligence: preliminary results and future directions," 2019, acesso em 14 maio 2025. [Online]. Available: <https://arxiv.org/pdf/1901.00246>
- [16] M. B. Cordeiro and W. M. N. Zola, "Knn paralelo em gpu para grandes volumes de dados com agregação de consultas," in *Anais do 24º Simpósio em Sistemas Computacionais de Alto Desempenho (WSCAD)*, 2023, pp. 253–264.
- [17] K. Arora, A. Raj, A. Goel, and S. Susan, "A hybrid model for combining neural image caption and k-nearest neighbor approach for image captioning," in *Soft Computing and Signal Processing*, ser. Advances in Intelligent Systems and Computing. Springer, 2022, vol. 1340, pp. 51–59.
- [18] A. M. Lourenço and A. E. Reis, "Aplicação da decisão multicritério ao problema de corte de estoque unidimensional," *Tendências em Matemática Aplicada e Computacional*, vol. 24, no. 1, pp. 117–135, 2023, acesso em 14 maio 2025. [Online]. Available: <https://www.scielo.br/j/tcam/a/ptJqJqQRfPH6Tcs46Tgzt6w/>
- [19] A. M. Fernandes, "Desenvolvimento da plataforma maut-i e implementação do aplicativo para cálculo da qualidade de uso de mídias interativas na primeira infância," Master's thesis, Universidade Federal dos Vales do Jequitinhonha e Mucuri, Diamantina, 2021, acesso em 14 maio 2025. [Online]. Available: https://sgppg.com.br/files/banco_de_teses/878.pdf
- [20] J. Pimentel-Filho *et al.*, "Modelagem de tópicos em documentos jurídicos via legal-bert," 2023, acesso em 14 maio 2025. [Online]. Available: <https://www.researchgate.net/publication/376807982>
- [21] C. M. T. Nascimento *et al.*, "Classificação interpretável e eficiente de decisões judiciais sobre habitação utilizando aprendizado baseado em subespaço e distância chordal adaptativa," in *Conferência Internacional Conjunta sobre Redes Neurais (IJCNN)*, Glasgow, 2020, pp. 1–8, acesso em 14 maio 2025. [Online]. Available: <https://ieeexplore.ieee.org/document/9152761>
- [22] S. Medhi and P. Tiwary, "Thermodynamics-inspired explanations of artificial intelligence," 2024, acesso em 14 maio 2025. [Online]. Available: <https://arxiv.org/pdf/2206.13475>
- [23] M. D. Kamalesh and B. Bharathi, "Personality prediction model for social media using machine learning technique," *Computers and Electrical Engineering*, vol. 100, p. 107852, 2022, acesso em 14 maio 2025. [Online]. Available: <https://www.sciencedirect.com/science/article/abs/pii/S0045790622001446>
- [24] G. Joshi *et al.*, "Explainable misinformation detection across multiple social media platforms," *IEEE Access*, vol. 11, pp. 23 634–23 646, 2023, acesso em 14 maio 2025.
- [25] Analytics Vidhya, "Entropy – a key concept for all data science beginners," 2020, acesso em 14 maio 2025. [Online]. Available: <https://www.analyticsvidhya.com/blog/2020/11/entropy-a-key-concept-for-all-data-science-beginners/>
- [26] S. A. Jesus Filho, "Análise de posts maliciosos na dark web utilizando aprendizado de máquina não supervisionado," 2023, acesso em 14 maio 2025. [Online]. Available: <https://repositorio.unesp.br/>
- [27] A. C. M. E. S. J. R. P. Lopes, "Otimização de aplicações de processamento de linguagem natural para análise de sentimentos," Master's thesis, Universidade Estadual Paulista, Instituto de Biociências, Letras e Ciências Exatas, São José do Rio Preto, 2021, acesso em 14 maio 2025. [Online]. Available: <https://repositorio.unesp.br/handle/11449/217830>
- [28] D. Moher, A. Liberati, J. Tetzlaff, D. G. Altman, and T. P. Group, "Preferred reporting items for systematic reviews and meta-analyses: The prisma statement," *PLoS Medicine*, vol. 6, no. 7, p. e1000097, 2009. [Online]. Available: <https://doi.org/10.1371/journal.pmed.1000097>