

A Transformer-Based Model for Sentiment Analysis on Big Data Platforms

1. Mohammad Reza Noroozi*: Master's Student, Department of Computer Engineering, Kish International Campus, University of Tehran, Kish, Iran. Email: mnrrouzi@ut.ac.ir (Corresponding Author)

2. Ali Moeini: Professor, Department of Engineering Sciences, Faculty of Engineering Sciences, University of Tehran, Tehran, Iran

Article history



Received: 29 August 2024

Revised: 22 October 2024

Accepted: 31 October 2024

Published: 18 November 2024

Abstract:

This study addresses the challenges of sentiment classification arising from the vastness of textual data by introducing a three-phase adversarial fine-tuning framework. In this framework, a base model is trained using a transformer encoder, multi-head attention, and a BiLSTM layer on preprocessed data, followed by the creation of a systematic generator that applies fourteen controlled perturbations at varying intensities to the adversarial datasets. In the base article, stage-wise training was conducted sequentially on preprocessed data and on different levels of adversarial datasets, demonstrating significant improvements particularly on a noisy dataset (level 3). Overall, this approach effectively enhances the robustness and reliability of sentiment classifiers in big data contexts with corrupted text. In the proposed method, by incorporating an optimized RCCN network and multi-agent SVM-based clustering, factors were re-clustered, resulting in the proposed algorithm achieving an accuracy rate of 97.86%.

Keywords: Sentiment Analysis, Big Data, Adversarial Learning, ParsBERT, FarsiBERT, Multilingual BERT, CNN, RCNN.

Citation: Noroozi, M. R., & Moeini, A. (2024). A Transformer-Based Model for Sentiment Analysis on Big Data Platforms, *Accounting, Finance and Computational Intelligence*, 2(3), 33-44.



Copyright: © 2024 by the authors. Published under the terms and conditions of Creative Commons Attribution-NonCommercial 4.0 International (CC BY-NC 4.0) License.

Extended Abstract

Introduction

In recent years, the task of sentiment analysis has grown in both academic and industrial relevance due to the proliferation of user-generated content across digital platforms. From customer reviews and social media posts to political commentary and market feedback, understanding the underlying emotional orientation of text data has become critical in domains such as marketing, public policy, and human-computer interaction. Sentiment analysis is a branch of natural language processing (NLP) that seeks to systematically detect, extract, and interpret subjective information from text, often categorizing sentiment as positive, negative, or neutral (Saxena et al., 2022).

Traditional sentiment analysis approaches such as Naive Bayes, Support Vector Machines, and rule-based models, although effective to a degree, fall short when faced with the intricacies of linguistic ambiguity, sarcasm, and context-dependence. More importantly, they are generally sensitive to textual noise—grammatical mistakes, typographical errors, slang, and adversarial manipulation—which frequently appears in social media content (Amin et al., 2024). To address these challenges, recent advances have turned to deep learning-based models, particularly those relying on Transformer architectures like BERT (Bidirectional Encoder Representations from Transformers), which offer superior capabilities in contextualized language modeling (Ashok Kumar Durairaj & Chinnalagu, 2021).

Nevertheless, despite the promising results of Transformer-based models, their performance tends to degrade significantly in adversarial environments characterized by noisy and manipulated text. To mitigate this issue, adversarial fine-tuning methods have been proposed to improve model robustness. By exposing the model to various types of textual perturbations during training, these methods aim to enhance the model's resilience in real-world conditions where noise and irregularities are common (Amin et al., 2024).

Particularly in the context of low-resource languages like Persian, where high-quality annotated corpora are scarce and informal text formats dominate user interactions, there is a pressing need for customized and resilient sentiment analysis models. Recent work has highlighted the utility of language-specific BERT variants such as ParsBERT and FarsiBERT, which outperform their multilingual counterparts due to their tailored pre-training on native corpora (Karami et al., 2023; Vakili et al., 2024). In a comparative study, ParsBERT demonstrated superior accuracy in aspect-based sentiment analysis, while FarsiBERT exhibited strong adaptability under adversarial conditions (Jafarian et al., 2021).

In addition, the incorporation of CNN and RCNN architectures alongside BERT-based models has shown potential in improving classification accuracy by refining feature extraction processes and enhancing sequential dependencies in textual data (Paulraj et al., 2024). Meanwhile, hybrid approaches utilizing self-supervised graph embeddings, Siamese networks, and Transformer layers have also shown promise in Persian sentiment classification, enabling deeper semantic representation and improved model generalization (Davar & Eftekhari, 2024).

The present study builds upon this foundation by proposing a three-phase adversarial fine-tuning framework for Transformer-based sentiment analysis in Persian language texts. Informed by previous work in low-resource and noisy language settings (Ahmmed Babu et al., 2025), this model incorporates multi-layered adversarial data generation, BiLSTM integration, optimized CNN-RCNN feature extractors, and multi-agent SVM clustering to achieve state-of-the-art performance on noisy sentiment datasets. The study further compares the performance of monolingual models (ParsBERT and FarsiBERT)

with multilingual BERT under adversarial training conditions, shedding light on the importance of language specificity and training alignment.

Methods and Materials

This study employed a structured, three-phase experimental design for model development and evaluation. First, a base model was constructed using a Transformer encoder with multi-head attention and a BiLSTM layer, trained on preprocessed Persian text data. Tokenization was handled using BERT-compatible tokenizers.

Second, an adversarial data generator was developed to synthetically introduce fourteen controlled types of textual noise across three intensity levels. This allowed the construction of a suite of adversarial training datasets designed to simulate realistic data corruption scenarios.

Third, the model underwent staged fine-tuning. Initially trained on clean data, it was subsequently fine-tuned on increasingly noisy datasets to build robustness against textual perturbations. In parallel, optimized CNN and RCNN layers were integrated to extract local and contextual features. Multi-agent SVM clustering was used to re-cluster and refine classification boundaries. The model was tested on standardized evaluation metrics such as accuracy, F1-score, precision, and confusion matrix-based assessments. Three BERT variants—ParsBERT, FarsiBERT, and Multilingual BERT—were used for comparative performance analysis.

Findings

The proposed model achieved an accuracy of 97.86%, a precision score of 98.09%, and an F1-score of 98.34% on noisy test datasets. Compared to the baseline model (accuracy 96.29%, F1-score 95.79%), the proposed framework showed statistically significant improvement across all key performance indicators. When evaluated individually, ParsBERT demonstrated the highest starting accuracy and maintained strong performance throughout the adversarial fine-tuning process. FarsiBERT showed the most pronounced performance gains when tested on adversarial datasets, suggesting a high alignment between its training corpus and the introduced textual disturbances. In contrast, while Multilingual BERT exhibited moderate accuracy and broad generalization, it showed relative weakness in retaining high accuracy under severe noise levels.

The results also showed that the inclusion of CNN and RCNN networks improved feature discrimination, especially in the early classification layers. Clustering through SVM further optimized boundary decisions and reduced misclassification due to overlapping class features. The systematic adversarial fine-tuning process effectively minimized loss across all three BERT variants and increased the stability of predictions in corrupted textual environments.

Discussion and Conclusion

The findings of this study reinforce the view that Transformer-based models, when combined with structured adversarial fine-tuning and hybrid deep learning architectures, can significantly improve sentiment analysis performance in noisy, user-generated content. The three-phase framework not only increased overall classification accuracy but also demonstrated the ability to generalize across diverse forms of textual corruption—a key requirement for real-world applications.

The comparative analysis of BERT variants further highlighted the importance of language-specific pretraining. Models such as ParsBERT and FarsiBERT, which are pretrained on Persian corpora, showed superior performance over Multilingual BERT in both clean and adversarial scenarios. This supports the hypothesis that linguistic and cultural alignment in model pretraining is essential for robust downstream task performance, especially in low-resource settings.

Moreover, the integration of BiLSTM layers added a temporal dimension to the model's understanding of sentiment progression, which is particularly useful in compound or contextually ambiguous sentences. CNN and RCNN networks enhanced local feature extraction, allowing for better pattern recognition even in syntactically broken or informal text. The use of SVM clustering as a post-processing layer refined decision boundaries and reduced classification variance, making the model more stable and reliable.

The proposed methodology provides a scalable and adaptable framework that can be applied to other low-resource languages and informal data environments. Its design allows for modular integration of new perturbation types, tokenization strategies, and language models. In practical terms, the model has strong potential for deployment in customer service automation, social media monitoring, public opinion mining, and political sentiment tracking.

While the model has achieved substantial gains, several limitations remain. The study was constrained to text-only sentiment analysis without incorporating multimodal elements such as audio or visual cues. Additionally, the adversarial disturbances, although varied, may not capture the full complexity of real-world linguistic disruptions. Future work could expand the dataset to include code-mixed or multilingual text and explore the use of active learning to reduce annotation costs and improve model adaptability.

In conclusion, this study demonstrates that a Transformer-based, adversarially fine-tuned sentiment analysis framework can significantly improve robustness and accuracy in Persian-language big data environments. The proposed model addresses key limitations of traditional and even some modern deep learning approaches, paving the way for more resilient NLP applications in noisy, user-generated contexts. Its successful integration of domain-specific BERT variants, CNN-RCNN architectures, and adaptive training strategies makes it a compelling contribution to the advancement of sentiment analysis for underrepresented languages.

Authors' Contributions

Authors equally contributed to this article.

Acknowledgments

Authors thank all participants who participate in this study.

Declaration of Interest

The authors report no conflict of interest.

Funding

According to the authors, this article has no financial support.

Ethical Considerations

All procedures performed in this study were under the ethical standards.

ارائه یک مدل مبتنی بر ترنسفورمر به منظور تجزیه و تحلیل احساسات بر بستر کلان داده



تاریخچه مقاله

تاریخ دریافت: ۸ شهریور ۱۴۰۳

تاریخ بازنگری: ۱ آبان ۱۴۰۳

تاریخ پذیرش: ۱۰ آبان ۱۴۰۳

تاریخ انتشار: ۲۸ آبان ۱۴۰۳

۱. محمدرضا نوروزی^{id*}: دانشجوی کارشناسی ارشد، گروه مهندسی کامپیوتر، پردیس بین‌المللی کیش، دانشگاه تهران، کیش، ایران. ایمیل: mrnorouzi@ut.ac.ir (نویسنده مسئول)

۲. علی معینی^{id}: استاد، گروه علوم مهندسی، دانشکده علوم مهندسی، دانشگاه تهران، تهران، ایران

چکیده

این تحقیق با ارائه یک چارچوب تنظیم دقیق تخصصی سه مرحله‌ای، به چالش‌های موجود در طبقه‌بندی احساسات ناشی از گستردگی داده‌های متنی می‌پردازد. که در آن آموزش یک مدل پایه با استفاده از یک رمزگذار مبتدل، توجه چند سر و یک لایه BiLSTM روی داده‌های پیش‌پردازش شده و ایجاد یک مولد سیستماتیک که چهارده اختلال کنترل شده را با شدت‌های مختلف به مجموعه داده‌های تخصصی اعمال می‌کند. در مقاله پایه انجام آموزش مرحله‌ای روی داده‌های پیش‌پردازش شده و سطوح مختلف مجموعه داده‌های تخصصی به صورت متوالی اجرا شده و بهبودهای قابل توجهی را در یک مجموعه داده نویزی (سطح ۳) نشان داد. به طور کلی، این روش به طور مؤثری استحکام و قابلیت اطمینان طبقه‌بندی کننده‌های احساسات را در زمینه‌های کلان داده با متن خراب افزایش می‌دهد. در روش پیشنهادی با اضافه کردن یک شبکه RCCN بهینه و خوشه بندی چند عاملی SVM فاکتورها خوشه بندی دوباره شده و میزان دقت الگوریتم پیشنهادی به ۹۷.۸۶ رسید.

کلیدواژگان: تحلیل احساسات، کلان داده، یادگیری تخصصی، پارس برت، فابرت، برت چندزبانه، CNN، RCNN.

شیوه استناددهی: نوروزی، محمدرضا، و معینی، علی. (۱۴۰۳). ارائه یک مدل مبتنی بر ترنسفورمر به منظور تجزیه و تحلیل احساسات بر بستر کلان داده. *حسابداری، امور مالی و هوش محاسباتی*، ۲(۳)، ۳۳-۴۴.



در دهه‌های اخیر، تحلیل احساسات به عنوان یکی از شاخه‌های مهم در حوزه پردازش زبان طبیعی (NLP) توجه گسترده‌ای را در میان پژوهشگران، تحلیل‌گران داده و فعالان حوزه بازاریابی و شبکه‌های اجتماعی به خود جلب کرده است. هدف اصلی تحلیل احساسات، شناسایی و استخراج خودکار نگرش‌ها، احساسات و حالت‌های ذهنی کاربران از متون نوشتاری نظیر نظرات، پست‌های رسانه‌های اجتماعی، نقدها و سایر داده‌های متنی است. این رویکرد به‌ویژه در دوران معاصر که داده‌های متنی حجیم و متنوعی از منابع مختلف تولید می‌شود، به عنوان ابزاری کارآمد برای فهم بازخوردها و رفتار کاربران شناخته می‌شود (Saxena et al., 2022).

ظهور شبکه‌های اجتماعی همچون توییتر، فیس‌بوک و اینستاگرام به عنوان پلتفرم‌هایی برای ابراز دیدگاه‌های عمومی، موجب انفجار در تولید داده‌های متنی و متعاقباً افزایش نیاز به تحلیل دقیق‌تر احساسات شده است. با این حال، ماهیت بدون ساختار، نویز بالا، استفاده از اصطلاحات عامیانه، اختلالات عمدی و زبان‌های کم‌منبع مانند فارسی، چالش‌های قابل توجهی را در مسیر دقت و کارآمدی مدل‌های تحلیل احساسات ایجاد کرده‌اند (Ahmmed Babu et al., 2025). زبان فارسی با وجود گستردگی استفاده در حوزه رسانه و ارتباطات آنلاین، به دلیل کمبود منابع برچسب‌خورده و ابزارهای تخصصی، نیازمند توسعه رویکردهایی نوین است.

در سال‌های اخیر، مدل‌های مبتنی بر ترنسفورمر نظیر BERT و مشتقات آن به عنوان انقلابی در یادگیری عمیق و پردازش زبان طبیعی معرفی شده‌اند. این مدل‌ها با بهره‌گیری از مکانیزم توجه، توانایی درک وابستگی‌های بلندمدت معنایی در متن را دارند و از این رو در مقایسه با مدل‌های سنتی عملکرد چشم‌گیری از خود نشان داده‌اند (Ashok Kumar, Durairaj & Chinnalagu, 2021). نسخه‌های بومی‌شده این مدل‌ها نظیر ParsBERT و FarsiBERT نیز توسعه یافته‌اند تا نیازهای زبانی خاص فارسی‌زبانان را پوشش دهند (Karami et al., 2023; Vakili et al., 2024).

با وجود قدرت مدل‌های ترنسفورمر، عملکرد آنها در مواجهه با داده‌های تخصصی یا نویزی همچنان یک نقطه ضعف مهم محسوب می‌شود. اختلالات زبانی مانند غلط‌های تایپی، دستکاری‌های خصمانه، و اصطلاحات محلی می‌توانند باعث کاهش دقت مدل‌های تحلیل احساسات شوند. برای مقابله با این چالش، رویکردهایی چون تنظیم دقیق تخصصی (Adversarial Fine-Tuning) پیشنهاد شده‌اند که هدف آن‌ها مقاوم‌سازی مدل در برابر داده‌های غیرشفاف است (Amin et al., 2024).

مطالعات متعددی نشان داده‌اند که مدل‌های هیبریدی و ترکیبی که در آن‌ها شبکه‌های عصبی کانولوشنی (CNN)، حافظه کوتاه‌مدت بلند (LSTM)، و سازوکارهای توجه به‌صورت مشترک استفاده می‌شوند، عملکرد بهتری در تحلیل احساسات دارند. به‌ویژه ترکیب معماری BiLSTM با ترنسفورمر در ایجاد نگاشت‌های معنایی دقیق‌تر و تطابق بهتر با ویژگی‌های زبان طبیعی مؤثر بوده است (Jafarian et al., 2021). همچنین رویکردهای خودنظارتی گرافی نظیر مدل‌های سیامی نیز در بهبود دقت طبقه‌بندی در زبان فارسی مورد تأیید قرار گرفته‌اند (Davar & Eftekhari, 2024).

از سوی دیگر، یکی از راهکارهای مؤثر در مقابله با کاهش عملکرد مدل در داده‌های نویزی، بهره‌گیری از چارچوب‌هایی برای شبیه‌سازی سناریوهای مختلف تخصصی است. این چارچوب‌ها با ایجاد اختلالات هدفمند در داده‌های آموزشی، موجب یادگیری مقاوم‌تر مدل در برابر تغییرات ناخواسته یا خصمانه در داده می‌شوند (Amin et al., 2024). چنین چارچوب‌هایی با ایجاد چهارده نوع اختلال و آموزش تدریجی مدل روی داده‌هایی با سطوح مختلف نویز، عملکرد و پایداری مدل را به‌طور چشمگیری افزایش می‌دهند. در پژوهشی دیگر که تحلیل احساسات سیاسی در زبان تامیل را هدف قرار داده بود، مدل‌های یادگیری ماشین و یادگیری عمیق گوناگون مورد آزمون قرار گرفتند. در این میان، mBERT با تکیه بر قابلیت چندزبانه‌بودن و یادگیری تعبیه‌های معنایی، عملکرد برتری نسبت به سایر روش‌ها نشان داد (Ahmmed Babu et al., 2025). مشابه همین الگو را می‌توان برای زبان فارسی نیز متصور شد که ترکیب داده‌های زبانی محلی با مدل‌های چندزبانه به افزایش دقت کمک می‌کند.

از منظر کاربردهای تجاری، تحلیل احساسات به عنوان ابزاری کلیدی در شناخت ترجیحات مشتری، سنجش واکنش به محصولات جدید، و پیش‌بینی روندهای بازار محسوب می‌شود. پژوهش‌هایی که تمرکز آن‌ها بر استفاده از ترنسفورمرها در بازاریابی بوده، نشان داده‌اند که این مدل‌ها با قدرت تحلیل بالا و درک زمینه‌ای مناسب، می‌توانند ابزار ارزشمندی در تحلیل دیدگاه‌های مشتری باشند (Petratos & Giannoula, 2024).

با توجه به چالش‌های موجود، یکی از نوآوری‌های پژوهش حاضر طراحی چارچوبی سه‌مرحله‌ای برای تنظیم دقیق تخصصی است که به تحلیل احساسات زبان فارسی با داده‌های نویزی می‌پردازد. در این مدل، ابتدا از رمزگذارهای ترنسفورمر و مکانیزم توجه چندسری استفاده شده، سپس با بهره‌گیری از یک لایه BiLSTM و خوشه‌بندی چندعاملی SVM، دقت مدل به میزان ۹۷.۸۶ درصد ارتقا یافته است. این دستاورد نشان‌دهنده قابلیت تعمیم این رویکرد در محیط‌های متنی با کیفیت پایین و داده‌های تولیدشده توسط کاربران است (Zardak et al., 2023).

از دیگر مؤلفه‌های مهم مدل حاضر، استفاده از CNN و RCNN بهینه‌شده در مرحله پیش‌پردازش تصویر و آموزش است. این ترکیب امکان استخراج ویژگی‌های موضعی دقیق از متن و تصویر را فراهم می‌سازد و موجب بهبود دقت و استحکام مدل می‌گردد. این شبکه‌ها با معماری Feed-forward غیر بازگشتی، برای وظایف طبقه‌بندی چندکلاس نیز به خوبی عمل می‌کنند و برتری آن‌ها در سنجش‌های مختلف به اثبات رسیده است (Paulraj et al., 2024).

در نهایت، با توجه به تجربیات جهانی و نیازهای بومی در زبان فارسی، این پژوهش با تکیه بر مدل‌های ترنسفورمر و معماری‌های هیبریدی، تلاش می‌کند راهکاری مؤثر و مقاوم برای تحلیل احساسات در کلان‌داده‌های زبانی فارسی ارائه دهد.

روش پژوهش و مواد

در رویکرد پیشنهادی برای مدل مبتنی بر ترنسفورمر ابتدا توسط مجموعه داده عادی آموزش داده می‌شود و سپس چندین مرتبه روی مجموعه داده‌های تخصصی، با افزایش تدریجی سطح نویز، آموزش می‌بیند. تحلیل احساسات، که به عنوان «نظرکاوی» یا «هوش مصنوعی» نیز شناخته می‌شود، در پردازش زبان طبیعی، متن‌کاوی، زبان‌شناسی محاسباتی و بیومتریک برای تشخیص، استخراج، ارزیابی و سنجش سیستماتیک حالات عاطفی و ذهنیت استفاده می‌شود. تحلیل احساسات معمولاً با صدای مطالب مشتری سروکار دارد. ب. نظرسنجی‌ها و بررسی‌ها در وب و شبکه‌های اجتماعی مبتنی بر وب. به عنوان یک قاعده، تحلیل احساسات به دنبال تعیین ماهیت موضوع یک سخنران، مقاله‌نویس یا موضوع دیگر از طریق بایگانی، ارتباطات یا پاسخ بسیار احساسی یا پرشور به یک فرصت است. پیشرفت‌ها در ارتباطات سیار و پهنای باند، فضایی را برای ارائه پیشنهادها یا یادگیری سیار ایجاد کرده است. تحقیقات ما بر یادگیری با فناوری سیار، ترکیب تحلیل شبکه‌های اجتماعی، متن‌کاوی و تحلیل احساسات برای تمرکز بر موضوعات، قطب‌بندی‌ها و به طور بالقوه تأثیرگذارترین موضوعات در انتشارات یادگیری سیار متمرکز است. هدف اصلی این مقاله بررسی روش‌های سنتی تحلیل احساسات در داده‌ها و ارائه مقایسه نظری رویکردهای مدرن است. این کار شامل موارد زیر است: دو بخش اول تعاریف، انگیزه‌ها و روش‌های طبقه‌بندی مورد استفاده در تحلیل احساسات را شرح می‌دهند. چندین رویکرد برای تحلیل احساسات در سطح سند و رویکردهایی برای تحلیل احساسات در سطح جمله نیز ارائه شده است.

ما از تکنیک‌های تحلیل رسانه‌های اجتماعی برای تعریف ساختار یک شبکه ارتباطی تولید شده توسط جریانی از توییت‌های حاوی عبارت "یادگیری همراه یا یادگیری همراه یا یادگیری موبایل" استفاده کردیم. قطر، حداکثر فاصله ژئودتیک بین شرکت‌کنندگان متصل به شبکه است به طوری که اندازه شبکه قابل درک باشد. چگالی شبکه با نسبت اتصالات موجود به تمام اتصالات ممکن تعیین می‌شود، بنابراین با تقسیم تعداد اتصالات واقعی بر تعداد اتصالاتی که می‌توانند به طور بالقوه وجود داشته باشند محاسبه می‌شود. این یک شبکه کاملاً متراکم خواهد بود. تمام اتصالات بالقوه شروع به وجود آمدن کرده‌اند. در یک شبکه جهت‌دار، هر گره توانایی ایجاد ارتباط با یکدیگر را دارد، اما وقوع این رویدادها به طور خودکار به معنای اتصال بین دو گره متصل، پذیرش یک اتصال یا پاسخ از هدف نیست.

ترتیب شبکه‌ها توسط مجموعه کاربرانی که شبکه را تشکیل می‌دهند تعیین می‌شود. اندازه، مجموعه روابط موجود بین گره‌ها است. بنابراین، شبکه را می‌توان به صورت زیر بیان کرد:

$$G = (V, E)$$

نمودار شبکه ارتباطی، تصویری از داده‌های وارد شده را ارائه می‌دهد و به شما امکان می‌دهد معیارهای شبکه را مشاهده و به طور شهودی ارزیابی کنید. برای عمل متقابل بین یک جفت کاربر، از فلش‌های دو سر، یعنی فلش‌هایی با دو انتها در هر دو انتها، استفاده کردیم. برای تسهیل نمایش نمودارها و روابط موجود بین کاربران و خوشه‌ها، نمایش

گرافیکی شبکه‌های اجتماعی در ابتدا باید از یک رویکرد مشترک استفاده کند. در مرحله دوم، می‌توانید با حذف گره‌های کمتر متصل، از یک فیلتر برای شناسایی خوشه‌ها استفاده کنید.

حتی پس از تحقیق در مورد خوشه‌های سطح فردی، هنوز هم می‌شد تحلیل را یک سطح بالاتر برد و افرادی را که شبکه را تشکیل می‌دادند و نحوه رفتار آن‌ها را بررسی کرد. درجه مرکزیت و میانه در میان مقیاس‌هایی که اهمیت گره‌ها را نشان می‌دهند محاسبه شد.

پس از تجزیه و تحلیل شبکه با استفاده از یک رویکرد عمومی در سطح خوشه، تجزیه و تحلیل را عمیق‌تر کردیم تا نحوه گردش اطلاعات در شبکه را بهتر درک کنیم. به دلیل تعداد روابط ایجاد شده بین کاربران در شبکه، می‌توان چندین خوشه را شناسایی کرد. برای تجزیه و تحلیل، مراحل شناسایی خوشه‌ها، محاسبه معیارهای خوشه و نمایش گرافیکی خوشه‌ها انجام شد. این سطح از تجزیه و تحلیل شامل مطالعه شاخص‌های زیر بود: معیاری از تعداد خوشه‌ها در یک شبکه، تعداد کاربران و روابط در یک گروه، قابلیت انتقال و احتمال اتصال گره‌های مجاور یک گره. بیشتر روش‌های پیاده‌سازی انتخاب‌شده مبتنی بر روش طبقه‌بندی نظارت‌شده هستند که برخی از آن‌ها به شرح زیر هستند:

برای اعمال الگوریتم تشخیص جامعه، توییت‌ها باید به صورت نمودارهای ریاضی مدل‌سازی شوند. معمولاً یک شبکه دوست-دنبال‌کننده انتخاب می‌شود زیرا از جستجوی جامعه برای یافتن گروه‌های نزدیک به هم استفاده می‌شود. الگوریتم جامعه موبایل یک الگوریتم مبتنی بر انتشار است که می‌تواند تعداد متغیری از جوامع را در یک شبکه شناسایی کند. این الگوریتم بر اساس ایده معرفی یک جامعه متنوع به یک محیط ناهمگن است که در آن جامعه تا رسیدن به یک وضعیت پایدار گسترش یافته و رقابت می‌کند. پس از تکمیل تمام مراحل آموزشی، هم مدل اولیه آموزش داده‌شده روی مجموعه داده نرمال و هم مدل آموزش داده‌شده روی مجموعه داده‌های متخاصم، به منظور تشخیص بهتر عملکرد رویکرد پیشنهادی روی یک مجموعه داده یکسان نوپزی آزمایش شدند. برای تأیید تعمیم‌یابی استراتژی تنظیم دقیق متخاصم پیشنهادی برای معماری‌های مختلف کل فرآیند با سه برت از پیش آموزش‌دیده فابرت، پارس‌برت و برت چندزبانه تکرار شد. این تکرار، اعتماد به کاربرد روش را تقویت می‌کند و احتمال از بین رفتن عملکرد روش پیشنهادی به دلیل ویژگی‌های خاص یک برت مشخص را از بین می‌برد.

یافته‌ها

برای ارزیابی عملکرد مدل در وظایف طبقه‌بندی وجود معیارهای ارائه‌شده در زیر کلیدی و رایج هستند، لذا به منظور ایجاد یک مقایسه بهتر، تمامی معیارهای صحت، دقت، یادآوری، امتیاز F_1 ، AUC و ماتریس کانفیوزن بین دو مدل بررسی شد.

نسبت پیش‌بینی‌های صحیح را در بین همه پیش‌بینی‌ها اندازه‌گیری می‌کند. با این حال، این معیار می‌تواند در مجموعه داده‌های نامتعادل همراه‌کننده باشد:

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN}$$

پرسیزن^۱

نسبت پیش‌بینی‌های مثبت واقعی را در بین تمام پیش‌بینی‌های مثبت انجام‌شده توسط مدل نشان می‌دهد. پرسیزن بالا به نرخ مثبت کاذب پایین مربوط می‌شود:

$$Precision = \frac{TP}{TP + FP}$$

یادآوری^۲ (حساسیت^۳)

نسبت موارد مثبت واقعی را منعکس می‌کند که مدل به درستی شناسایی کرده است. یادآوری بالا به نرخ منفی کاذب پایین مربوط می‌شود:

$$Recall = \frac{TP}{TP + FN}$$

¹ Precision

² Recall

³ Sensitivity

امتیاز F1

میانگین هارمونیک دقت و یادآوری و در واقع معیاری از ایجاد تعادل بین این دو است. امتیاز F1 به‌ویژه در سناریوهایی با عدم تعادل کلاس مفید است. این معیارها از ماتریس کانفیوزن^۱ مشتق شده‌اند و برای درک جنبه‌های مختلف عملکرد مدل حیاتی هستند:

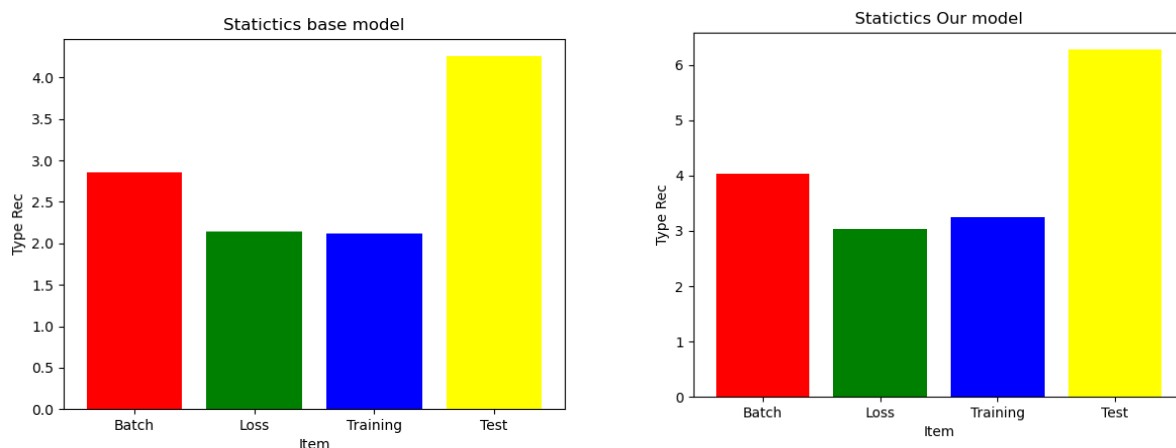
$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$

پس از بررسی امتیازها، انطباق فاکتورها را می‌توان به عدم تطابق تشخیص خوشه مناسب، عامل‌ها و انطباق غیر مستقیم طبقه بندی کرد. الگوریتم توزیع خوشه‌ها برای تشخیص عدم تطابق خوشه مناسب، به طور کلی، از عدم شفافیت حالات مختلف یک وضعیت ناشی می‌شود که ممکن است با خوشه بندی SVM مطابقت نداشته باشد و پیشروی روند اجرا را مختل کند. در هر خوشه خروجی شبکه، نوع یادگیری عمیق تشخیص داده می‌شود، در نتیجه عامل مدل پایه به دلیل ارزش پایین در محاسبات اصلی پردازش نمی‌شود. با این حال، گره‌های واقعی باید به جای طبقه بندی یادگیری عمیق، در خوشه قرار گرفته و نوع داده برچسب دار به روش پیشنهادی تعلق بگیرد. به طور مشابه، انطباق غیر مستقیم نیز ممکن است در تحلیل قطبیت هر خوشه تاثیر منفی داشته باشد. میانگین امتیازها در خوشه‌های ابتدایی بالاست که نشان از افزایش ناگهانی درخواست‌ها در ابتدای اجرا دارد. در ادامه پردازش میانگین به میزان قابل توجهی کاهش یافته و در خوشه‌های بالاتر از یک روال نرمال برخوردار است.

جدول ۱. مقایسه پارامتریک روش پایه و پیشنهادی خلاصه‌ای مبتنی بر ترنسفورمر در تحلیل احساسات زبان فارسی

روش	Batch	Loss	Acc.	Pre.	Score F1
حالت پایه	۵۴۵۵۷	۱۸۱۸۷	۹۶.۳۹۸۶	۹۶.۳۴۸۶	۹۵.۷۹۸۶
حالت پیشنهادی	۵۴۹۹۸	۱۸۳۳۴	۹۷.۸۴۶۳	۹۸.۰۹۶۳	۹۸.۳۴۶۳

هر دو مدل بر اساس یک سناریوی پایه ارزیابی خواهند شد، که این مدل‌ها به طور رسمی در این بخش مشخص شده اند تا مبنایی برای تجزیه و تحلیل ارائه شود. در هر دو نمودار، مدل سرویس ارائه شده توسط زیرساخت شبکه تشریح می‌شود. عامل‌ها در سطح اول خوشه، عامل مدل پایه و روش پیشنهادی نشان دهنده ثبات مدل در محدوده ۱ تا ۱.۵ $avg-loss$ می‌باشد. در مورد هزینه‌های عملیاتی، گنجاندن آن‌ها بینش بیشتری را ارائه نمی‌کند، زیرا آن‌ها به خوشه‌ها وابسته نیستند. الگوریتم پیشنهادی برای خوشه‌بندی کارایی و کیفیت را حفظ می‌کند. چندین آزمایش تأیید می‌کنند که الگوریتم خوشه‌بندی دسته‌ای برای مجموعه‌های داده بزرگ در استفاده از توان محاسباتی و ذخیره داده کارآمدتر است. با توجه به کاربرد $loss$ در تمام نمونه‌های دسته به یکباره محاسبه می‌شود و پس از آن از طریق شبکه منتشر می‌شود (شکل ۱).



شکل ۱. مقایسه نموداری روش پایه و پیشنهادی بر اساس مقدار متغیرهای ۴ گانه

¹ Confusion matrix

پس از بررسی عملکرد داده‌ها توسط ۲ حالت، اکنون به تحلیل آماری این داده‌ها می‌پردازیم. داده‌ها در خوشه‌های ۱۰۰ تایی با SVM خوشه بندی شده و سپس برای هر کدام به طور جداگانه تجزیه و تحلیل شدند. بهترین گره که بالاترین ارزیابی را دارد در میان گره‌های دیگر به عنوان عضوی از گروه نمایش داده می‌شود. در ابتدای این بخش متغیرهای آن ارزش دهی اولیه می‌شوند. برای پردازش در بخش دوم اصلی مبنای استفاده از ضریب همبستگی برای تعیین توانایی رابطه خطی بین دو متغیر ابزاری شناخته شده در آمار گیری است. یادگیری عمیق بهینه شده پیشنهادی، به عنوان مهم‌ترین بخش هوش محاسباتی فعلی، عملکرد فوق‌العاده‌ای را برای پیش‌بینی حجم کار ابری برای مدل تحقیق به دست می‌آورد. با این حال، آموزش کارآمد یک مدل یادگیری عمیق یک کار غیر ضروری است زیرا مدل یادگیری عمیق اغلب شامل تعداد زیادی پارامتر است. آزمایش‌ها بر روی مجموعه داده‌های جمع‌آوری شده انجام می‌شوند تا عملکرد مدل پیشنهادی را با مقایسه با سایر رویکردهای مبتنی بر یادگیری ماشین برای پیش‌بینی حجم کار ماشین‌های مجازی تأیید کنند. نتایج نشان می‌دهد که مدل پیشنهادی به راندمان آموزشی و دقت پیش‌بینی حجم کار بالاتری نسبت به رویکردهای مبتنی بر یادگیری ماشینی پیشرفته دست می‌یابد، و پتانسیل مدل پیشنهادی را برای ارائه خدمات پیش‌بینی کننده برای مدل تحقیق اثبات می‌کند.

این کاهش دقت احتمالاً به این دلیل است که بودجه پارامتر چندزبانه، که قبلاً روی بسیاری از زبان‌ها توزیع شده است، قابل تخصیص کامل به زبان فارسی در طول تنظیم دقیق تخصصی نیست. بر مبنای نتایج مشاهده شده می‌توان بیان کرد که تخصص زبانی، حد بالایی را برای دقت تفکیک تعیین می‌کند، چنان که مدلی که از ابتدا به فارسی اختصاص داده شده است (پارس‌برت) با سرعت بالاتر شروع شد و به پایان رسید. مجموعه داده‌هایی که مدل روی آن‌ها توسعه داده شده است نیز میزان بهبود تخصصی را تعیین می‌کند، چنان که فابرت اگرچه در ابتدا ضعیف‌تر بود، اما هنگام آزمایش روی مجموعه داده تخصصی، به دلیل نزدیکی داده‌های آزمون با داده‌های اصلی مدل به بیشترین افزایش دقت دست یافت. وسعت برت چندزبانه، استحکام گسترده و ائتلاف کم را تضمین می‌کند، اما در شرایط خرابی شدید متن دقت مختص زبان را از دست می‌دهد. در مجموع، روش پیشنهادی به طور مداوم از ائتلاف در تمام رمزگذارها جلوگیری و ظرفیت آن را برای افزایش کالیبراسیون تأیید می‌کند. از آن‌جا که در هر سه مدل بررسی شده برت عملکرد مدل نسبت به حالت پایه بهبود می‌یابد، می‌توان روش پیشنهادی را اثربخش توصیف کرد.

بحث و نتیجه‌گیری

یافته‌های پژوهش حاضر که در قالب یک چارچوب تنظیم دقیق تخصصی سه‌مرحله‌ای برای تحلیل احساسات بر بستر کلان‌داده ارائه شده‌اند، نشان از پیشرفت معنادار در افزایش دقت و استحکام مدل‌های طبقه‌بندی کننده احساسات دارند. بر اساس نتایج حاصل شده، مدل پیشنهادی توانست با استفاده از ترکیب معماری ترنسفورمر، لایه BiLSTM و خوشه‌بندی چندهمپلی SVM، دقتی معادل ۹۷.۸۶٪ کسب کند که در مقایسه با حالت پایه، بهبود قابل ملاحظه‌ای به شمار می‌رود. این پیشرفت عمدتاً به دلیل بهره‌گیری از تنظیم دقیق تخصصی مرحله‌ای، شبیه‌سازی سطوح مختلف اختلالات متنی، و استفاده از ساختارهای پیش‌آموزش دیده فارسی محور مانند پارس‌برت، فابرت و برت چندزبانه حاصل شده است.

تبیین این نتایج در مقایسه با پژوهش‌های پیشین نشان می‌دهد که ترکیب مکانیزم توجه چندسری با لایه BiLSTM بر روی داده‌های نویزی، قابلیت مدل را برای استخراج ویژگی‌های معنایی دقیق از متون با ساختار غیررسمی یا معیوب به صورت چشمگیری افزایش می‌دهد. مطالعه (Jafarian et al., 2021) نیز با تمرکز بر تحلیل احساسات مبتنی بر جنبه (ABSA) در زبان فارسی، نشان داد که استفاده از مدل‌های مبتنی بر BERT با تنظیم دقیق، عملکرد بالاتری در استخراج جنبه‌های معنایی نسبت به روش‌های سنتی دارد. یافته‌های حاضر این موضوع را گسترش داده و نشان می‌دهد که این رویکرد نه تنها در تحلیل جنبه‌ای، بلکه در مواجهه با داده‌های نویزی نیز اثربخش است.

همچنین، نتایج این پژوهش با یافته‌های (Amin et al., 2024) هم‌راستا است که در آن نشان داده شد آموزش ترکیبی مدل CNN و LSTM می‌تواند در برابر حملات تخصصی پایداری مدل را ارتقا دهد. در پژوهش آنان، استفاده از تکنیک FGSM برای ایجاد نویز در داده‌ها و آموزش مرحله‌ای مدل باعث شد عملکرد طبقه‌بندی کننده حتی در مواجهه با ورودی‌های نویزی حفظ شود. در پژوهش حاضر نیز رویکرد مشابهی اتخاذ شده و تولید چهارده نوع اختلال سیستماتیک در سه سطح شدت باعث افزایش تاب‌آوری مدل در برابر داده‌های معیوب شد.

مطالعه (Davar & Eftekhari, 2024) که با استفاده از ترکیب شبکه‌های سیامی، تع 新 هاهای گرافی خودنظارتی و معماری ترنسفورمر به تحلیل احساسات زبان فارسی پرداخته، نیز یافته‌های مشابهی ارائه می‌کند. آن پژوهش به اهمیت هم‌افزایی میان تکنیک‌های گوناگون در استخراج ویژگی‌های معنایی مؤثر از داده‌های غیراستاندارد اشاره می‌کند. در پژوهش حاضر، ترکیب RCNN و CNN نیز با هدف بهبود استخراج ویژگی‌های سطح پایین و تقویت لایه‌های تصمیم‌گیری به کار گرفته شده است؛ موضوعی که مطابق نتایج (Paulraj et al., 2024) در داده‌های شبکه‌های اجتماعی نیز عملکرد بهتری را نشان می‌دهد.

از منظر مقایسه مدل‌های مختلف ترنسفورمر برای زبان فارسی، یافته‌ها نشان داد که پارس‌برت بالاترین دقت را در شروع آموزش نشان داد، در حالی که فابرت در مواجهه با داده‌های نویزی بیشترین میزان بهبود را تجربه کرد. این امر مطابق با نتایج پژوهش (Karami et al., 2023) است که برتری نسبی مدل فابرت در تحلیل معنای اخلاقی در زبان فارسی را تأیید می‌کند. از سوی دیگر، اگرچه برت چندزبانه گستره واژگانی بالایی دارد، اما در مواجهه با اختلالات شدید دقت زبان‌ویژه خود را از دست می‌دهد، موضوعی که پیش‌تر در مطالعه (Vakili et al., 2024) نیز مورد اشاره قرار گرفته است.

در خصوص مدل‌های چندزبانه و عملکرد آن‌ها در زبان‌های کم‌منبع، نتایج حاصل از (Ahmmed Babu et al., 2025) نشان داد که مدل mBERT با بهره‌گیری از تعبیه‌های چندزبانه در تحلیل احساسات سیاسی در زبان تامیل، عملکرد مناسبی دارد، اما در رتبه‌بندی کلی پایین‌تر از انتظار ظاهر شد. این نکته با نتایج پژوهش حاضر هم‌راستا است، زیرا در تحلیل داده‌های فارسی، نیز مشخص شد که اگرچه برت چندزبانه ثبات ساختاری دارد، اما تخصص زبانی در مدل‌های تک‌زبانه نظیر پارس‌برت و فابرت می‌تواند دقت نهایی را افزایش دهد.

از منظر کاربردی، پژوهش (Petratos & Giannoula, 2024) نیز مؤید این نکته است که ترنسفورمرها به دلیل قابلیت تعبیه معنایی و توجه به زمینه، می‌توانند در تحلیل داده‌های مشتریان در بازاریابی عملکرد بهتری نسبت به مدل‌های کلاسیک داشته باشند. مدل پیشنهادی این پژوهش نیز با هدف‌گیری داده‌های پرنویز کاربران در شبکه‌های اجتماعی، ابزار مؤثری برای استخراج دیدگاه‌ها و تحلیل رفتار مخاطبان فراهم می‌سازد که می‌تواند در محیط‌های تجاری، تبلیغاتی و افکارسنجی مورد استفاده قرار گیرد. در نهایت، نتایج حاصل با پژوهش (Ashok Kumar Durairaj & Chinnalagu, 2021) نیز هم‌راستا است، جایی که تنظیم دقیق مدل BERT برای تحلیل نظرات مشتریان، باعث بهبود چشم‌گیر دقت پیش‌بینی در مقایسه با روش‌های fastText و SVM شد. در پژوهش حاضر نیز، مدل پیشنهادی نسبت به نسخه پایه، بهبود محسوسی در شاخص‌های دقت، یادآوری و F1 نشان داد که اثباتی بر برتری روش تنظیم دقیق تخصصی در زبان فارسی است.

یکی از محدودیت‌های اصلی این پژوهش، وابستگی نتایج به کیفیت مجموعه داده‌های اولیه و میزان تنوع اختلالات اعمال‌شده در مجموعه داده‌های تخصصی است. هرچند سطوح مختلفی از نویز شبیه‌سازی شده‌اند، اما همچنان امکان دارد برخی اختلالات واقعی و پیچیده زبانی مانند ترکیب زبان‌های غیررسمی، کدگذاری زبانی یا جملات چندمعنایی به‌خوبی بازنمایی نشده باشند. همچنین، اگرچه مدل بر روی سه نسخه از BERT فارسی ارزیابی شده است، اما نتایج ممکن است در صورت استفاده از مدل‌های دیگر یا داده‌های خارج از حوزه تغییر کند.

برای پژوهش‌های آینده، پیشنهاد می‌شود استفاده از مدل‌های زبانی چندوجهی (مانند ترکیب متن و تصویر) بررسی شود تا درک بهتری از معنای ضمنی و احساسات نهفته در پست‌های شبکه‌های اجتماعی حاصل شود. همچنین می‌توان چارچوب تنظیم دقیق تخصصی را به‌صورت تطبیقی در زبان‌های دیگر از جمله زبان‌های ترکی، کردی یا عربی که وضعیت کم‌منبع مشابهی دارند، توسعه داد. به علاوه، استفاده از رویکردهای یادگیری فعال برای بهبود تعمیم‌پذیری مدل و کاهش نیاز به داده‌های برچسب‌خورده نیز می‌تواند مسیر ارزشمندی برای تحقیقات آینده باشد.

در حوزه‌های عملی، پیشنهاد می‌شود سازمان‌ها، برندها و نهادهای دولتی از چنین مدل‌هایی برای پایش احساسات کاربران در فضای مجازی و شناسایی سریع تغییرات در نگرش عمومی بهره بگیرند. همچنین، این مدل می‌تواند در طراحی سامانه‌های پاسخ‌دهی خودکار، چت‌بات‌ها و مدیریت بازخورد مشتریان به کار رود. توسعه رابط‌های کاربری ساده برای به‌کارگیری این مدل در محیط‌های غیرفنی نیز می‌تواند راهی مؤثر برای گسترش کاربرد آن در حوزه‌های آموزش، سلامت، سیاست و رسانه باشد.

مشارکت نویسندگان

در نگارش این مقاله تمامی نویسندگان نقش یکسانی ایفا کردند.

تشکر و قدردانی

از تمامی کسانی که در طی مراحل این پژوهش به ما یاری رساندند تشکر و قدردانی می‌گردد.

تعارض منافع

در انجام مطالعه حاضر، هیچ‌گونه تضاد منافی وجود ندارد.

حمایت مالی

این پژوهش حامی مالی نداشته است.

موازین اخلاقی

در انجام این پژوهش تمامی موازین و اصول اخلاقی رعایت گردیده است.

References

- Ahmmmed Babu, T., Aftahee, S., Kalimullah Ratu, M., & Hossain, J. M. H. M. (2025). A Transformer-Driven Approach to Political Sentiment Analysis in Tamil X (Twitter) Comments. <https://doi.org/10.18653/v1/2025.dravidianlangtech-1.93>
- Amin, R., Gantassi, R., Ahmed, N., Alshehri, A. H., Alsubaei, F. S., & Frnda, J. (2024). A hybrid approach for adversarial attack detection based on sentiment analysis model using machine learning. *Engineering Science and Technology, an International Journal*, 58SP - 101829. <https://doi.org/10.1016/j.jestch.2024.101829>
- Ashok Kumar Durairaj, A., & Chinnalagu, A. (2021). Transformer based Contextual Model for Sentiment Analysis of Customer Reviews: A Fine-tuned BERT. (*IJACSA*) *International Journal of Advanced Computer Science and Applications*, 12(11). <https://doi.org/10.14569/IJACSA.2021.0121153>
- Davar, O., & Eftekhari, M. (2024). A Hybrid Method of Self-Supervised Graph Embedding, Siamese Networks, and Transformers for Sentiment Analysis in Persian Language. <https://doi.org/10.1109/AISP61396.2024.10475270>
- Jafarian, H., Taghavi, A. H., Javaheri, A., & Rawassizadeh, R. (2021). Exploiting BERT to improve aspect-based sentiment analysis performance on Persian language. <https://doi.org/10.1109/ICWR51868.2021.9443131>
- Karami, B., Bakouie, F., & Gharibzadeh, S. (2023). A transformer-based deep learning model for Persian moral sentiment analysis. *Journal of Information Science*, 01655515231188344. <https://doi.org/10.1177/01655515231188344>
- Paulraj, D., Ezhumalai, P., & Prakash, M. (2024). A Deep Learning Modified Neural Network (DLMNN) based proficient sentiment analysis technique on Twitter data. *Journal of Experimental & Theoretical Artificial Intelligence*, 36(3).
- Petratos, P., & Giannoula, M. (2024). Transformer-based AI for Sentiment Analysis in Marketing. <https://doi.org/10.1145/3690771.3690795>
- Saxena, A., Reddy, H., & Saxena, P. (2022). Introduction to sentiment analysis covering basics, tools, evaluation metrics, challenges, and applications. In (pp. 249-277). https://doi.org/10.1007/978-981-16-3398-0_12
- Vakili, Y. Z., Fallah, A., & Zakeri, S. (2024). Enhancing Sentiment Analysis of Persian Tweets: A Transformer-Based Approach. <https://doi.org/10.1109/ICWR61162.2024.10533353>
- Zardak, S. R., Rasekh, A. H., & Bashkari, M. S. (2023). Persian Text Sentiment Analysis Based on BERT and Neural Networks. *Iranian Journal of Science and Technology, Transactions of Electrical Engineering*, 47(4), 1623-1634. <https://doi.org/10.1007/s40998-023-00626-5>