



Semnan University

Journal of Modeling in Engineering

Journal homepage: <https://modelling.semnan.ac.ir/>

ISSN: 2783-2538



Research Article

Aspect-based sentiment analysis based on users' comments in an online marketplace

Fatemeh Moodi¹ , Amir Jahangard-Rafsanjani^{2,*} , Fatemeh Sadri³

^{1,2,3} Department of Computer Engineering, Yazd University, Yazd, Iran

PAPER INFO

Paper history:

Received:

Revised:

Accepted:

Keywords:

Sentiment analysis;

FastText;

Different aspects of comments;

Deep learning;

Convolutional neural network;

BERT.

ABSTRACT

In recent years, with the expansion of online shopping and the importance of user feedback in improving the quality of products and services, sentiment analysis of user reviews has become one of the most important tools and opportunities for online businesses and services. In this research, various approaches based on Convolutional Neural Networks (CNN), BERT language model, and FastText language model have been examined for sentiment analysis of user reviews about mobile phones in aspects such as camera, battery, and price in an online marketplace. In this regard, the data were labeled based on aspects and categorized into three sentiments: positive, negative, and neutral. Additionally, to improve the performance of the proposed approaches and address data imbalance, data augmentation methods were utilized and their impact was analyzed. Finally, using the proposed models, very high accuracy in detecting positive, negative, and neutral reviews in each aspect was achieved. According to the results, the proposed CNN-based approach performed better than the other two proposed methods on the given data.

DOI: <https://doi.org/>

© 2024 Published by Semnan University Press.

This is an open access article under the CC-BY 4.0 license. (<https://creativecommons.org/licenses/by/4.0/>)

* Corresponding author.

E-mail address: jahangard@yazd.ac.ir

How to cite this article:

تحلیل احساسات مبتنی بر جنبه روی نظریات کاربران در یک بازارگاه برخط

فاطمه مودی^۱، امیر جهانگرد-رفسنجانی^{۲*} و فاطمه صدری^۳

اطلاعات مقاله	چکیده
دریافت مقاله:	در سال‌های اخیر، با توجه به گسترش خریدهای اینترنتی و اهمیت بازخورد کاربران در بهبود کیفیت محصولات و خدمات، تحلیل احساسات روی نظریات کاربران به یکی از مهم‌ترین ابزارها و فرصت‌های پیش‌روی کسب‌وکارها و خدمات برخط تبدیل شده است. در این پژوهش رویکردهای مختلف مبتنی بر شبکه عصبی پیچشی، مدل زبانی برت و همچنین مدل زبانی فست‌تکست برای تحلیل احساسات نظریات کاربران در مورد تلفن همراه در جنبه‌های دوربین، باتری و قیمت در یک بازارگاه اینترنتی بررسی شده است. در همین راستا، داده‌ها مبتنی بر جنبه و بر اساس سه احساس مثبت، منفی و خنثی برچسب‌گذاری شد. همچنین برای بهبود عملکرد رویکردهای پیشنهادی و عدم توازن بین داده‌ها، از روش‌های افزایش داده‌ها استفاده و تأثیر آن‌ها بررسی شده است. در نهایت با استفاده از مدل‌های پیشنهادی، دقت بسیار بالایی در تشخیص نظریات مثبت، منفی و خنثی در هر جنبه به‌دست آمد. بر اساس نتایج، رویکرد پیشنهادی مبتنی بر شبکه عصبی پیچشی، عملکرد بهتری نسبت به دو روش پیشنهادی دیگر روی داده‌های مورد نظر داشته است.
واژگان کلیدی:	
تحلیل احساسات، فست‌تکست، جنبه‌های مختلف نظرها، یادگیری عمیق، شبکه عصبی پیچشی، برت.	
DOI: https://doi.org/	
© 2024 Published by Semnan University Press.	
This is an open access article under the CC-BY 4.0 license. (https://creativecommons.org/licenses/by/4.0/)	

۱-مقدمه

با گسترش استفاده از فضای وب، خریدهای اینترنتی و تعاملات برخط کاربران در سال‌های اخیر، بازخورد و نظریات کاربران به‌عنوان منبعی ارزشمند از محتوای مصرفی در دنیای وب معرفی می‌شود. این عوامل، به سازمان‌ها امکان می‌دهد تا با تجزیه و تحلیل دقیق‌تر این ورودی‌ها و نظریات کاربران، خدمات و محصولات خود را بهبود بخشند. همچنین، از این اطلاعات برای بهبود تجربه کاربری، بهبود عملکرد کلی سازمان و بهره‌برداری بهتر از منابع خود استفاده کنند [۱]. برای دستیابی به این اهداف،

سازمان‌ها از تحلیل احساسات بهره می‌برند. این فرآیند شامل ارزیابی و تجزیه و تحلیل نظریات کاربران برای شناخت و فهم بهتر احساسات، نگرانی‌ها و تمایلات آن‌ها است. با بهره‌گیری از این ابزارهای تحلیلی، سازمان‌ها می‌توانند به شناخت عمیق‌تری از واکنش‌ها و نیازهای مشتریان خود دست پیدا کنند و در نتیجه تصمیم‌گیری‌های بهتری داشته باشند.

امروزه، با پیشرفت تکنولوژی و افزایش خدمات برخط، تحلیل احساسات نظریات مشتریان در یک بازارگاه برخط

استناد به این مقاله: نحوه استناد فارسی در اینجا درج گردد.

* پست الکترونیک نویسنده مسئول: jahangard@yazd.ac.ir

۱. دانشکده مهندسی برق و کامپیوتر، گروه مهندسی کامپیوتر، دانشگاه یزد، یزد، ایران.

(CNN)، مدل زبانی برت^۲ (BERT) با تنظیم اختصاصی روی داده‌های پژوهش و مدل زبانی فست‌تکست^۴ هستند. هدف این پژوهش، بهبود دقت و کارایی این مدل‌ها در تحلیل احساسات فارسی است. برای این منظور، مجموعه داده‌ای شامل نظرات کاربران درباره تلفن همراه، با برچسب‌های مثبت، منفی و خنثی بر اساس جنبه‌های مختلف ایجاد شده است. عملکرد مدل‌ها برای هر جنبه به صورت جداگانه ارزیابی شده است. همچنین، با بررسی و آزمایش روش‌های افزایش داده‌ها^۵ و متوازن‌سازی داده‌ها^۶، تلاش شده است تا دقت مدل‌ها بهبود یابد. در نهایت، با پیاده‌سازی شبکه‌های پیشنهادی، تنظیم دقیق فرآیندها و انتخاب مقادیر بهینه، عملکرد این مدل‌ها روی داده‌های موردنظر ارزیابی شده است.

۲- مبانی نظری و پیشینه تحقیق

رسانه‌های اجتماعی در سال‌های اخیر رشد چشمگیری داشته‌اند. با افزایش روزافزون استفاده از اینترنت، وب به یکی از مهم‌ترین منابع اطلاعاتی تبدیل شده است [۳۹]. میلیون‌ها نفر نظرها و احساسات خود را در فروشگاه‌های اینترنتی، وبلاگ‌ها و شبکه‌های اجتماعی به اشتراک می‌گذارند. حجم گسترده‌ی محتوای تولیدشده توسط کاربران؛ بسترهای اینترنتی و رسانه‌های اجتماعی را به یکی از گسترده‌ترین منابع داده‌ی افکار عمومی تبدیل کرده است. از این رو تحلیل این داده‌ها برای نظارت بر افکار عمومی و کمک به تصمیم‌گیری یک ضرورت است. تحلیل احساسات یکی از روش‌هایی است که برای تحلیل داده‌های تولیدشده توسط کاربران استفاده می‌شود [۵ و ۶].

۲-۱- سطوح تحلیل احساسات

تحلیل احساسات معمولاً در سه سطح سند^۷، جمله^۸ و جنبه^۹ بررسی می‌شود. در سطح سند به بررسی این که کل سند حاوی نظر مثبت یا منفی است پرداخته می‌شود و دسته‌بندی بر اساس احساس و نظر کلی سند صورت می‌گیرد [۷]. سند را می‌توان مجموعه‌ای از جمله‌ها در نظر گرفت و طبقه‌بندی در سطح سند زمانی بهترین عملکرد را دارد که سند توسط یک شخص و یا در رابطه با موضوعی خاص نوشته شده باشد [۸]. در سطح جمله تمرکز بر جمله

اهمیت بسیار زیادی دارد و می‌تواند به پیشرفت آن کسب-وکار و افزایش رضایت مشتریان کمک کند. از جمله دلایل اهمیت تحلیل احساسات در کسب‌وکارها؛ درک نیازهای مشتری، کیفیت محصولات و خدمات، افزایش رضایت مشتریان، جذب مشتریان جدید و پیش‌بینی مشکلات آینده است.

کلمات، ابزاری برای بیان افکار، نظرهای و نیازهای افراد هستند. در عین حال، یک کلمه ممکن است در همه موارد معنی یکسانی نداشته باشد و معنی آن وابسته به کاربرد آن شکل می‌گیرد. این پیچیدگی در زبان‌ها از مواردی است که تحلیل احساسات را با چالش مواجه می‌کند. تحلیل احساسات در زبان فارسی، با وجود پیشرفت‌های قابل توجه در سال‌های اخیر، همچنان با چالش‌هایی روبه‌رو است [۲]. از مهم‌ترین این چالش‌ها می‌توان به موارد زیر اشاره کرد: (۱) کمبود داده و منابع: در مقایسه با زبان‌هایی مانند انگلیسی، داده‌های کمتری برای تحلیل احساسات فارسی در دسترس است. این کمبود می‌تواند فرایند آموزش و تنظیم مدل‌ها را محدود کند [۳].

(۲) پیچیدگی ساختاری زبان فارسی: زبان فارسی دارای ساختار زبانی پیچیده‌ای است که شامل ترکیب‌های خاص کلمات و اصطلاحات متنوع می‌شود. همچنین، وجود اشتباه‌های املائی و نویزهای نوشتاری، دقت مدل‌های تحلیل احساسات را کاهش می‌دهد.

(۳) چندمعنایی واژگان: بسیاری از واژگان فارسی بسته به زمینه، معانی متفاوتی دارند. درک صحیح این واژه‌ها نیازمند تحلیل دقیق متن و شناخت درست کاربرد آن‌هاست [۲]. برای بهبود دقت و کارایی تحلیل احساسات در زبان فارسی، توسعه مدل‌های پیشرفته و راهکارهایی برای رفع این چالش‌ها ضروری است [۴].

این پژوهش با هدف تحلیل احساسات نظرات کاربران درباره جنبه‌های مختلف تلفن همراه در یک بازارگاه برخط و بررسی فرآیند پیش‌پردازش داده‌ها در زبان فارسی انجام شده است. در این راستا، رویکردهای مختلف تحلیل احساسات بررسی شده و رویکردی مناسب برای تحلیل احساسات و پیش‌پردازش داده‌های فارسی ارائه گردیده است. رویکردهای پیشنهادی شامل شبکه عصبی پیچشی^۲

⁶ Data Balancing

⁷ Document Level

⁸ Sentence Level

⁹ Aspect Level

² Convolutional Neural Network

³ Bidirectional Encoder Representations from Transformers

⁴ FastText

⁵ Data Augmentation

حوزه‌های مختلف می‌توانند معانی و مفاهیم متعددی داشته باشند. بنابراین ممکن است یک کلمه مثبت در حوزه‌ای خاص، در حوزه‌ای دیگر مفهوم متفاوتی داشته باشد [۱۳].

۲-۳- رویکردهای تحلیل احساسات مبتنی بر یادگیری ماشین

رویکردهای یادگیری ماشین برای طبقه‌بندی قطبیت^{۱۶} احساسات (برای مثال مثبت، منفی یا خنثی بودن)، به روش‌های یادگیری نظارت‌شده^{۱۷}، یادگیری بدون نظارت، یادگیری نیمه‌نظارتی^{۱۸} و یادگیری تقویتی^{۱۹} تقسیم می‌شود [۵]. یادگیری نظارت‌شده زمانی استفاده می‌شود که طبقه‌بندی دارای مجموعه مشخصی از کلاس‌ها باشد و نیازمند داده‌های آموزشی دارای برچسب است [۱۴]. روش‌های بدون نظارت زمانی مناسب است که برچسب‌گذاری برای مجموعه داده مقدور نباشد [۱۵]. روش نیمه‌نظارتی را برای مجموعه داده‌های بدون برچسب که شامل برخی از نمونه‌های برچسب‌دار است، می‌توان استفاده کرد. یادگیری تقویتی از مکانیسم‌های آزمون و خطا برای کمک به عامل، در تعامل با محیط اطراف، استفاده می‌کند. عامل، بازخوردهایی از محیط می‌گیرد و تجربه‌هایی از محیط کسب می‌کند و از آن‌ها برای به‌دست‌آوردن بیشترین پاداش استفاده می‌کند [۵]. استفاده از رویکردهای یادگیری ماشین معمولاً نتایج خوبی را به همراه دارد، ولی برای دستیابی به عملکرد خوب به مجموعه داده‌های آموزشی بزرگ نیاز دارند.

۲-۴- رویکردهای ترکیبی

رویکرد ترکیبی، رویکردهای مبتنی بر واژگان و یادگیری ماشین را ترکیب می‌کند. علت اصلی استفاده از رویکرد ترکیبی، بهره‌بردن از دقت بالای رویکردهای یادگیری ماشین و ثبات رویکردهای مبتنی بر واژگان و بهره‌گیری از مزایای آن‌ها است [۵ و ۱۶].

۲-۵- یادگیری عمیق در تحلیل احساسات

اصطلاح یادگیری عمیق، به شبکه‌های عصبی با لایه‌های متعدد پرسپترون اشاره دارد که الهام گرفته از شبکه عصبی مغز انسان است. با این معماری، مدل‌های پیچیده‌تری را

است و هدف تعیین مثبت، منفی یا خنثی بودن جمله‌ی مورد نظر است [۹]. یکی از معایب تحلیل احساسات در سطح سند این است که تشخیص احساسات متفاوت در مورد موجودیت‌های مختلف دشوار است. در این مواقع تحلیل احساسات در سطح جمله می‌تواند عملکرد بهتری داشته باشد. سطح جنبه، به بررسی جنبه‌ها یا ویژگی‌های گوناگون موجودیت‌ها و تحلیل احساسات مربوط به هر کدام از آن‌ها می‌پردازد و در نتیجه تحلیل دقیق‌تری ارائه می‌دهد [۱۰]. در این سطح از تحلیل احساسات؛ ابتدا موضوعات و موجودیت‌ها و جنبه‌های بیان‌شده در متن شناسایی و استخراج می‌شود، سپس به تحلیل احساس برای هر کدام از جنبه‌ها پرداخته می‌شود. این رویکرد با شناسایی جنبه‌ها و موضوعات بیان‌شده توسط نویسنده، اطلاعات کاملی را در اختیار قرار می‌دهد. در واقع در این سطح مشخص می‌شود که افراد دقیقاً چه چیزی را دوست دارند و چه چیزی را دوست ندارند [۱۱].

در اکثر مواقع، رویکردهای تحلیل احساسات به سه دسته‌ی کلی شامل رویکردهای مبتنی بر واژگان^{۱۰}، رویکردهای یادگیری ماشین^{۱۱} و رویکردهای ترکیبی^{۱۲} تقسیم می‌شود. یادگیری ماشین پرکاربردترین رویکرد در تحلیل احساسات است و برای طبقه‌بندی احساسات به الگوریتم‌های یادگیری ماشین و ویژگی‌های زبان متکی است [۵]. در ادامه به بررسی هر کدام از این رویکردها می‌پردازیم.

۲-۲- رویکردهای تحلیل احساسات مبتنی بر واژگان

در رویکردهای مبتنی بر واژگان، که به آن رویکرد مبتنی بر دانش^{۱۳} نیز می‌گویند، تحلیل احساسات مبتنی بر احساسات واژگان و مجموعه‌ای از واژگان با بار احساسی شناخته‌شده و ازپیش تعیین‌شده است. در این روش از شباهت‌های عاطفی کلمات و عبارات استفاده می‌شود و دقت آن به وزن‌های ازقبل‌آموزش‌دیده^{۱۴} وابسته است [۱۲ و ۱۳]. رویکرد مبتنی بر واژگان در تحلیل احساسات در سطح جمله بسیار کاربردی است. این رویکرد نیازی به داده‌های آموزشی ندارد و از این‌رو، آن را به‌عنوان رویکرد بدون نظارت^{۱۵} می‌توان در نظر گرفت. از سوی دیگر، مشکل اصلی این رویکرد وابستگی به حوزه‌ی کاربردی است، زیرا کلمات در

¹⁵ Unsupervised

¹⁶ Polarity

¹⁷ Supervised Learning

¹⁸ Semi-Supervised Learning

¹⁹ Reinforcement Learning

¹⁰ Lexicon-Based

¹¹ Machine Learning

¹² Hybrid

¹³ Knowledge-Based

¹⁴ Pre-Learned

زبان به یک سری از پیش نیازهای مشخص، وابسته است و استفاده‌ی مستقیم از روش‌ها، ابزارها و منابع توسعه‌یافته برای زبان انگلیسی در فارسی با محدودیت‌هایی مواجه است. بنابراین، تحلیل احساسات در زبان فارسی به دلیل دارا بودن ویژگی‌های خاص این زبان، نسبت به زبان انگلیسی، نیازمند اتخاذ روش‌ها و مدل‌های منحصر به فرد است [۲]. یکی از انواع این مدل‌ها، که مبتنی بر ارزیابی وابستگی بین توالی‌های ورودی/خروجی است، مدل ترنسفورمر نام دارد. این مدل به رمزگذار و رمزگشا وابسته است. رمزگذار دنباله ورودی را می‌گیرد و آن را به برداری با ابعاد بالاتر نگاشت می‌کند. سپس این بردار توسط رمزگشا به دنباله‌ای از خروجی نگاشت می‌شود. از جمله مدل‌سازی‌های زبان از-پیش آموزش دیده^{۲۵} مبتنی بر مدل ترنسفورمر، مدل‌های برت [۲۵] و GPT [۲۲] است و به عنوان مدل‌های پرکاربرد و مؤثر در حوزه مدل‌سازی زبان شناخته می‌شوند [۲۳]. برای زبان فارسی نیز چندین مدل جاسازی کلمات^{۲۷} از جمله GloVe، Word2Vec و فست‌تکست ارائه شده است [۲۴].

۲-۶-۱- مدل زبانی برت

برت یک چارچوب یادگیری ماشین منبع باز و یک مدل زبانی قدرتمند برای پردازش زبان طبیعی است. برت، یک مدل از پیش آموزش دیده و دوطرفه و مبتنی بر ترنسفورمر است که در سال ۲۰۱۸ معرفی شد [۲۵]. در واقع برت یک مدل یادگیری عمیق است که در آن هر عنصر خروجی به هر عنصر ورودی متصل می‌شود. وزن‌های بین این عناصر، به صورت پویا و بر اساس اتصال آن‌ها محاسبه می‌شود [۲۶]. مدل زبان آموزشی دوطرفه در استخراج ویژگی‌ها و تفسیر متن، بهتر از مدل زبان یک طرفه عمل می‌کند. در واقع در معماری دو طرفه متن به صورت دوسویه، از اول تا آخر و از آخر به اول، پردازش می‌شود. برای دستیابی به نمایش برداری بهتر، مدل برت از ساختار ترنسفورمر برای یادگیری اطلاعات زمینه‌ای کلمات ورودی استفاده می‌کند. توانایی تجزیه و تحلیل متن به صورت دوطرفه، جلورو^{۲۸} و عقب‌رو^{۲۹}، به این نوع مدل‌ها اجازه می‌دهد تا دانش قبلی خود را به کار خاصی اختصاص دهند؛ دانشی که از طریق یک مرحله

روی مجموعه داده‌ای بسیار بزرگ می‌توان آموزش داد و نتایج بهتری تولید کرد [۵]. در سال‌های اخیر، تحلیل احساسات مبتنی بر روش‌های یادگیری عمیق به عنوان روشی قدرتمند برای تحلیل احساسات دقیق مطرح شده است [۱۷]. شبکه عصبی بازگشتی^{۲۰}، شبکه عصبی پیچشی و حافظه طولانی کوتاه مدت^{۲۱}، رایج‌ترین مدل‌های یادگیری عمیق در تحلیل احساسات هستند [۱۸].

۲-۵-۱- شبکه عصبی پیچشی

شبکه عصبی پیچشی نوع خاصی از شبکه عصبی پیشخور^{۲۲} است که در ابتدا در حوزه بینایی کامپیوتری به کار می‌رفت؛ اما امروزه در زمینه‌های مختلف مانند سیستم‌های توصیه گر و پردازش زبان طبیعی^{۲۳} به نتایج موفقی دست یافته است [۵ و ۱۹]. شبکه عصبی پیچشی شامل یک لایه ورودی، یک لایه خروجی و یک لایه پنهان است. در معماری این شبکه، لایه پنهان شامل چندین لایه پیچشی، لایه ادغام^{۲۴} و لایه کاملاً متصل^{۲۵} است [۲۰]. لایه ورودی وظیفه دریافت داده و انتقال آن به لایه‌های بعدی را دارد. لایه پیچشی برای تشخیص الگوها و استخراج ویژگی‌ها در داده‌های ورودی طراحی شده است و محاسبات شبکه عصبی در این لایه انجام می‌شود. اندازه فضای ویژگی که در لایه پیچشی به دست آمده، توسط لایه ادغام کاهش می‌یابد و اطلاعات مهم حفظ می‌شود. پس از استخراج ویژگی‌ها از طریق لایه‌های پیچشی و ادغام، از لایه‌های کاملاً متصل به عنوان لایه‌های پایانی برای اتصال ویژگی‌های استخراج شده و تصمیم‌گیری نهایی، با توجه به اطلاعات به دست آمده، استفاده می‌شود. برای مثال در مسائل دسته‌بندی چندکلاسه، با استفاده از یک تابع فعال‌ساز، احتمال مربوط به هر کلاس محاسبه می‌شود. در نهایت در لایه خروجی، نتایج نهایی مدل به دست می‌آید [۱۹ و ۲۱].

۲-۶-۲- تحلیل احساسات مبتنی بر مدل‌های زبانی

در سال‌های اخیر تحقیقات گسترده‌ای برای ساخت مدل‌های مختلف تحلیل احساسات برای زبان انگلیسی، فارسی و سایر زبان‌ها با استفاده از الگوریتم‌های یادگیری ماشین و یادگیری عمیق انجام شده و مدل‌سازی زبان محبوبیت پیدا کرده است [۱۷]. تحلیل احساسات در هر

²⁵ Fully Connected

²⁶ Pre-Trained

²⁷ Word Embeddings

²⁸ Forward

²⁹ Backward

²⁰ Recurrent Neural Network (RNN)

²¹ Long Short-Term Memory (LSTM)

²² Feedforward Neural Network

²³ Natural Language Processing

²⁴ Pooling

را می‌دهند. در فست‌تکست برای ساخت بردار، ابتدا تمام ویژگی‌های تعبیه‌شده مربوط به کلمات در متن جمع‌آوری می‌شوند. سپس با میانگین‌گیری از تمام این ویژگی‌ها، یک بردار نهایی برای متن به‌دست می‌آید.

- آموزش مدل: با استفاده از بردارهای نمایش متن و برجسب‌های موجود، مدل برای طبقه‌بندی آموزش داده می‌شود. این مدل می‌تواند از الگوریتم‌هایی مانند شبکه عصبی پیچشی استفاده کند.
- به‌طور خلاصه، فست‌تکست روشی قدرتمند برای طبقه‌بندی متن است و از ویژگی‌های جاسازی و ترکیب کلمات برای دسته‌بندی دقیق و سریع متن‌ها در کلاس‌های مختلف استفاده می‌کند [۲۹ و ۳۰].

۲-۷- پیشینه تحقیق

تحلیل احساسات یکی از مسائل مهم در زمینه علوم داده و پردازش زبان طبیعی است که در سال‌های اخیر توجه محققان، صنعت و جوامع علمی را به خود جلب کرده است و تأثیر مهمی در تصمیم‌گیری‌ها، توسعه محصولات، بازاریابی، ابزارهای پشتیبانی از مشتری و حتی سیاست‌گذاری دارد.

جوړک و همکاران [۳۲] الگوریتمی جدید برای تحلیل احساسات مبتنی بر واژگان ارائه کردند که تمرکز اصلی آن تحلیل محتوای توییت‌هاست. این الگوریتم، از دو جزء کلیدی شامل نرمال‌سازی احساسات و تابع ترکیبی مبتنی بر شواهد^{۳۱} تشکیل شده است. از این الگوریتم برای تخمین شدت احساسات به جای برجسب مثبت و منفی و طبقه‌بندی احساسات استفاده شده است و از مجموعه‌ای داده‌ای شامل ۱/۶ میلیون توییت استفاده کردند.

سینگلا و همکاران [۳۳] از الگوریتم‌های یادگیری ماشین برای طبقه‌بندی احساسات استفاده کردند. آن‌ها، داده‌های جمع‌آوری شده از وب‌سایت آمازون را با استفاده از روش‌های نیویز^{۳۲} و ماشین‌بردار پشتیبان^{۳۳} به نظرهای مثبت و منفی دسته‌بندی کردند. در این تحقیق، روش ماشین‌بردار-پشتیبان با دقت ۸۱/۷۵ درصد بهترین عملکرد را داشت. ابراهیمی و همکاران [۳۴] روشی نظارت‌شده برای توسعه‌ی مجموعه‌ی واژگان با بار احساسی مشخص با استفاده از استخراج ویژگی‌ها برای زبان فارسی پیشنهاد کردند. در این

طولانی از قبل آموزش‌دیده، در مجموعه‌های عظیم به‌دست آمده است. در واقع، شبکه را با استفاده از دو روش استخراج ویژگی و تنظیم دقیق می‌توان اختصاصی کرد.

۲-۶-۲- مدل زبانی فست‌تکست

مدل فست‌تکست [۲۸] روشی پیشرفته در پردازش زبان طبیعی برای طبقه‌بندی متن و توسعه‌یافته توسط فیس‌بوک است. این روش برای پردازش متن‌های طولانی، با تعداد زیاد کلاس یا برجسب، بسیار کارآمد است. در این روش، مفاهیم بسته‌ی کلمات^{۳۰} و جاسازی کلمات ترکیب می‌شود و مزایایی مانند سرعت بالا و دقت معقول را ارائه می‌دهد [۲۹]. در ادامه مراحل اصلی فست‌تکست برای طبقه‌بندی متن آمده است.

- ساخت لغت‌نامه و جاسازی کلمات: مدل فست‌تکست از الگوریتم Word2Vec برای ایجاد جاسازی کلمات استفاده می‌کند. در این مرحله، هر کلمه به برداری عددی تبدیل می‌شود که ویژگی‌های مختلف زبان را مشخص می‌کند. این نمایش برداری، بر اساس معانی مشترک کلمات در میان متن ایجاد می‌شود.
- ساخت n-gram: در فست‌تکست علاوه بر تک‌واژه‌ها، از n-gram نیز استفاده می‌شود. n-gram مجموعه‌ای از n کلمه‌ی متوالی در متن است. این مفهوم به تحلیل و فهم متن‌ها، بر اساس رابطه و ترتیب عناصر متوالی و درک ارتباطات بین کلمات در چارچوبی خاص از متن کمک می‌کند. استفاده از n-gram در پردازش زبان طبیعی بسیار مفید است. به‌عنوان مثال، در حوزه‌ی متن‌کاوی و طبقه‌بندی متن، n-gram در زمینه‌ی تشخیص الگوها و روابط بین کلمات و توکن‌ها و همچنین تحلیل دقیق‌تر متن کمک‌کننده است.
- ساخت بردارهای نمایش متن: ساخت بردارهای نمایش متن گامی اساسی در بسیاری از حوزه‌های کاربردی پردازش زبان طبیعی، از جمله طبقه‌بندی متن با استفاده از روش‌هایی مانند فست‌تکست است. این فرآیند شامل تبدیل محتوای متنی به بردارهای عددی است که الگوریتم‌های یادگیری ماشین قادر به درک و پردازش آن‌ها هستند. این بردارها، ویژگی‌های اساسی متن را به‌گونه‌ای ثبت می‌کنند که به مدل امکان پیش‌بینی یا انجام وظایف مختلف پردازش زبان طبیعی

³² Naïve Bayes

³³ Support Vector Machine

³⁰ Bag of Words (BoW)

³¹ Evidence-Based

با استفاده از روش‌های پردازش زبان طبیعی و الگوریتم‌های مختلف، به شناسایی الگوهای احساسی می‌پردازند. طبقه‌بندی احساسات در سطح سند موضوعی اساسی در تجزیه و تحلیل احساسات است و برای درک محتوای تولیدشده توسط کاربران در شبکه‌های اجتماعی یا بررسی محصول بسیار مهم است.

تانگ و همکاران [۴۰] با مطرح کردن تحلیل احساسات در سطح سند به عنوان یک چالش، یک مدل شبکه عصبی برای یادگیری اسناد متوالی، برای طبقه‌بندی احساسات معرفی کردند. روش آن‌ها بر اساس اصل ترکیب‌بندی است که بیان می‌کند معنای عبارتی طولانی مثل یک جمله یا یک سند، به معانی اجزای تشکیل‌دهنده آن بستگی دارد. آن‌ها از رویکردهای شبکه عصبی پیچشی، حافظه طولانی کوتاه‌مدت و شبکه عصبی بازگشتی استفاده کردند. در این پژوهش، طبقه‌بندی احساسات در سطح سند روی چهار مجموعه داده در مقیاس بزرگ انجام شد. نتایج تجربی نشان می‌دهد که مدل پیشنهادی عملکرد بهتری نسبت به چندین الگوریتم دیگر داشت. تحلیل احساسات در سطح جمله یکی از سطوح اصلی تحلیل احساسات است.

تحقیقات زیادی در این زمینه، برای شناسایی قطبیت جملات، بر اساس نشانه‌های زبانی استخراج شده از محتوای متن صورت گرفته است. یکی از چالش‌های مطرح در این سطح این است که انواع مختلف جملات می‌توانند احساسات متفاوتی را بیان کنند. به عنوان مثال، در جمله خبری "خوب است"، قطبیت احساس مثبت است. در عین حال در جمله پرسشی "آیا خوب است؟" قطبیت احساس مبهم است.

چن و همکاران [۹] یک روش تقسیم‌و‌غلبه^{۳۶} ارائه کردند که ابتدا جملات را به انواع مختلف طبقه‌بندی می‌کنند؛ سپس تحلیل احساسات را به طور جداگانه روی جملات هر نوع انجام می‌دهند. آن‌ها با بررسی رابطه بین تعداد اهداف نظری بیان‌شده و احساسات بیان‌شده در یک جمله، چارچوبی جدید برای بهبود تحلیل احساسات از طریق طبقه‌بندی نوع جمله پیشنهاد کردند. اهداف نظری، به هر موجودیت یا جنبه‌ای از جمله اشاره دارد که نظری در مورد آن ابراز شده است. آن‌ها در کار خود از مدلی مبتنی بر شبکه عصبی برای طبقه‌بندی جملات به سه نوع، با توجه به تعداد اهداف

پژوهش، از نظرهای جمع‌آوری‌شده از یک سایت فروش آنلاین به عنوان مجموعه‌ای داده استفاده شد. این مجموعه شامل ۷,۵۰۰ نظر در زمینه‌ی دوربین دیجیتال، لپ تاپ، تلویزیون، تبلت و تلفن همراه است که به صورت دستی جمع‌آوری شده و از آن برای انتخاب ویژگی و ایجاد لیستی از کلمات احساسی استفاده شده است. برای ارزیابی عملکرد رویکرد پیشنهادی، از روش ماشین‌بردار پشتیبان استفاده شد و دقت ۸۰ درصد به دست آمد.

ارمنی و همکارانش [۳۵] رویکردی ترکیبی مبتنی بر روش ماشین‌بردار پشتیبان و جنگل تصادفی پیشنهاد کردند. داده‌های این پژوهش از وبسایت آمازون جمع‌آوری شد و نتایج آن نشان داد که رویکرد ترکیبی با دقت ۸۳/۴ عملکرد بهتری از سایر الگوریتم‌ها به صورت جداگانه داشت.

کیم و همکاران [۳۶] شبکه عصبی پیچشی را برای طبقه‌بندی احساسات در کار خود پیشنهاد کردند. آن‌ها از طریق آزمایش‌هایی با سه مجموعه داده‌ی شناخته‌شده، نشان دادند که استفاده از لایه‌های پیچشی متوالی برای متن‌های نسبتاً طولانی مؤثرتر است و عملکرد بهتری از سایر مدل‌های یادگیری عمیق دارد.

سانی و همکاران [۳۷] معماری جدیدی به نام TextConvoNet و مبتنی بر شبکه عصبی پیچشی برای مسائل طبقه‌بندی متن دوکلاسه و چندکلاسه ارائه کردند. آن‌ها با استفاده از این مدل ویژگی‌های درون جمله‌ای و بین جمله‌ای را با n-gram از داده‌های ورودی استخراج می‌کنند. نتایج تجربی حاصل از این پژوهش نشان می‌دهد که مدل ارائه‌شده نسبت به سایر مدل‌های مورد استفاده برای طبقه‌بندی متن، عملکرد بهتری دارد.

جانگ و همکاران [۳۸] مدلی ترکیبی از حافظه کوتاه‌مدت-دوطرفه^{۳۴} و شبکه عصبی پیچشی مبتنی بر توجه^{۳۵} را پیشنهاد کردند. برای این کار از مجموعه داده‌ی مربوط به نظرهای فیلم‌ها در IMDB^{۳۶} استفاده شد. نتایج حاصل از مدل ترکیبی پیشنهادی، با دقت ۸۷/۳۳ درصد، نشان‌دهنده‌ی بهبود نتایج نسبت به مدل‌های حافظه-طولانی کوتاه‌مدت و شبکه عصبی پیچشی است.

همانطور که قبلاً گفته شد؛ تحلیل احساسات در سطح‌های مختلف از جمله سطح سند، جمله و جنبه مورد بررسی قرار می‌گیرد. این سطوح تحلیل احساسات در جملات و سندها،

³⁶ Internet Movie Database

³⁷ Divide and Conquer

³⁴ Bi-LSTM

³⁵ Attention-Based

فارسی مبتنی بر ترکیب روابط معنایی و جنگل تصادفی پرداختند. در این پژوهش، نظرات کاربران سایت دیجی کالا مورد بررسی قرار گرفته است. ابتدا عملیات پاکسازی داده‌ها صورت گرفته و سپس ویژگی‌ها بر اساس روابط معنایی فست‌تکست استخراج شده است. در مرحله بعد، عملیات کاهش ویژگی توسط شبکه‌های یادگیری عمیق انجام شده و در نهایت، برای طبقه‌بندی نظرات از الگوریتم جنگل تصادفی استفاده شده است. نتایج نشان‌دهنده دقت ۹۸/۵ درصد، صحت ۹۷ درصد، فراخوان ۹۸ درصد و معیار F1 برابر با ۹۷ درصد می‌باشد [۴۳].

قاسمپور مرزبالی از تحلیل احساسات مبتنی بر جنبه با استفاده از برت استفاده کرد. این مطالعه به تحلیل احساسات ریزدانه (ABSA) پرداخته است. برای مقابله با محدودیت اندازه مجموعه داده‌های آموزشی، از تولید متن مبتنی بر برت همراه با الگوریتم‌های فیلتر متن استفاده شده است. این رویکرد به بهبود کیفیت مجموعه داده‌ها و افزایش عملکرد مدل‌ها در طبقه‌بندی احساسات سطح جنبه منجر شده است [۴۴].

فائزه فروتن و محمد ربیعی با هدف بهبود استخراج ویژگی‌ها از نظرات فارسی و افزایش دقت تحلیل احساسات، چارچوبی جدید برای تحلیل احساسات در سطح جمله ارائه می‌دهند. این چارچوب مبتنی بر مدل‌های برت، CNN-BiLSTM و طبقه‌بندی‌کننده XGBoost است. نتایج نشان‌دهنده دقت ۹۳/۷۴ درصد در طبقه‌بندی احساسات متون فارسی است که نشان می‌دهد ترکیب مدل‌های مذکور می‌تواند به بهبود دقت تحلیل احساسات در زبان فارسی منجر شود [۴۵].

لی و همکاران به تحلیل لحن و احساسات نظرات شبکه‌های اجتماعی در پلتفرم‌های رسانه‌ای جدید با استفاده از مدل برت پرداختند. در این مطالعه، مدل برت به دلیل قابلیت‌های رمزگذاری دوطرفه و درک متنی‌اش، برای تحلیل احساسات و لحن نظرات در پلتفرم‌های رسانه‌ای جدید انتخاب شده است. از طریق فرآیند پیش‌آموزش و تنظیم دقیق، نگرش‌ها و قطبیت‌های احساسی نظرات به همراه ویژگی‌های لحنی مانند شادی، خشم و طعنه به دقت شناسایی شدند. تحلیل دقیق لحن و احساسات معنایی نظرات در پلتفرم‌های رسانه‌ای جدید، درک عمیقی از ترجیحات کاربران و روندهای افکار عمومی فراهم می‌کند که می‌تواند به بهینه‌سازی توصیه‌های محتوا، بهبود

ظاهر شده در یک جمله، استفاده کردند. سپس هر گروه از جملات به طور جداگانه، برای طبقه‌بندی احساسات، به عنوان ورودی شبکه عصبی پیچشی در نظر گرفته شد. نتایج تجربی حاصل از این پژوهش نشان می‌دهد که طبقه‌بندی نوع جمله می‌تواند عملکرد تحلیل احساسات در سطح جمله را بهبود بخشد. همچنین رویکرد پیشنهادی به نتایج مطلوبی در چندین مجموعه داده معیار رسید.

در سطح جنبه، یافتن احساسات با توجه به جنبه‌های خاص موجودیت‌ها صورت می‌گیرد. از آنجا که ممکن است در یک متن، احساسی واحد نسبت به موضوعی منتقل نشود، نیاز به شناخت دقیق تر نظرها و احساسات، که به عنوان تحلیل احساسات مبتنی بر جنبه شناخته می‌شود، در سال‌های گذشته مورد توجه زیادی قرار گرفته است [۴۱].

الگریوتی و همکاران [۴۲] با هدف کمک به نهادهای دولتی برای کسب دیدگاه جامع تر از نگرش افراد، یک رویکرد ترکیبی تحلیل احساسات مبتنی بر جنبه را پیشنهاد کردند که دامنه واژگان و قوانین را برای تجزیه و تحلیل نظرهای برنامه‌های هوشمند، یکپارچه می‌کند. هدف مدل پیشنهادی استخراج جنبه‌های مهم از این نظرها و طبقه‌بندی احساسات مربوط به آن‌ها است. با توجه به نتایج گزارش شده، دقت استخراج جنبه با در نظر گرفتن جنبه‌های ضمنی به طور قابل توجهی بهبود یافته است. همچنین رویکرد پیشنهادی، در طبقه‌بندی احساسات، از روش‌های یادگیری ماشین همچون ماشین بردار پشتیبان بهتر عمل کرد.

ماجومدر و همکاران [۴۱] تحلیل احساسات سطح جنبه را به صورت دو بخش متمایز استخراج جنبه و برچسب‌گذاری جنبه‌ها با قطبیت احساسات در نظر گرفتند. با وجود اینکه این دو کار متمایز هستند، اما ارتباط زیادی با هم دارند. رویکرد پیشنهادی این فرضیه را مطرح می‌کند که انتقال دانش از یک مدل استخراج جنبه از پیش‌آموزش دیده، به عملکرد مدل‌ها در طبقه‌بندی قطبیت احساسات می‌تواند کمک کند. آن‌ها بر اساس این فرضیه، جاسازی کلمات را در طول استخراج جنبه به دست آورند؛ سپس آن را به عنوان ورودی مدل‌های طبقه‌بندی قطبیت احساسات در نظر گرفتند. نتایج تجربی نشان داد که اطلاعات اضافه شده، به طور قابل توجهی عملکرد مدل‌های پایه در طبقه‌بندی احساسات را بهبود بخشیده است.

روستائی و همکاران به تحلیل احساسات کاربران در زبان

تجربیات کاربری و افزایش کارایی عملیاتی این پلتفرم‌ها کمک کند. نتایج نشان می‌دهد که مدل برت قادر است به‌طور دقیق نگرش‌های احساسی و ویژگی‌های لحنی نظرات را شناسایی کند [۴۶].

طلعت چهار مدل عمیق مبتنی بر ترکیب برت با الگوریتم‌های BiLSTM و BiGRU را پیشنهاد می‌دهد. این مدل‌ها بر روی سه مجموعه داده آموزش داده شدند و نتایج نشان داد که ترکیب برت با BiLSTM و BiGRU می‌تواند عملکرد بهتری نسبت به استفاده تنها از برت داشته باشد [۴۷].

مانوئل-ایلی و همکاران به بررسی جامع الگوریتم تحلیل احساسات با استفاده از مدل برت پرداخته و مبانی نظری،

معیارهای ارزیابی، یافته‌ها و تحلیل مقایسه‌ای را ارائه می‌دهد. نویسندگان با استفاده از مدل برت، الگوریتمی برای تحلیل احساسات توسعه داده‌اند. این الگوریتم قادر است با دقت بالایی احساسات موجود در متون را شناسایی کند و در کاربردهای مختلفی مانند پیش‌بینی قطبیت احساسات و پیش‌بینی احساسات مورد استفاده قرار گیرد. این مطالعه نشان می‌دهد که استفاده از مدل‌های مبتنی بر یادگیری عمیق، به‌ویژه برت، می‌تواند بهبود قابل‌توجهی در دقت و کارایی سیستم‌های تحلیل احساسات ایجاد کند [۴۸].

در جدول ۱، مقایسه پژوهش‌های تحلیل احساسات در سال‌های اخیر آمده است.

جدول ۱- مقایسه پژوهش‌های تحلیل احساسات

پژوهشگر/سال	روش تحلیل احساسات	داده‌های مورد استفاده	الگوریتم‌های یادگیری ماشین	دقت مدل
جانگ و همکاران [۳۸] (۲۰۲۰)	یادگیری عمیق ترکیبی	داده‌ی مربوط به نظرهای فیلم‌ها در IMDB	مدلی ترکیبی از حافظه کوتاه-مدت دوطرفه و شبکه عصبی پیچشی مبتنی بر توجه	دقت ۸۷/۳۳ درصد
روستائی و همکاران [۴۳] (۱۴۰۱)	ترکیب روابط معنایی و جنگل تصادفی	نظرهای سایت دیجی کالا	جنگل تصادفی	۹۸/۵ درصد
قاسمپور مرزبالی [۴۴] (۱۴۰۳)	ABSA مبتنی بر برت	مجموعه داده‌های ABSA	برت و تولید متن	بهبود عملکرد مدل‌ها
فائزه فروتن و محمد [۴۵] (ربیع ۲۰۲۳)	برت، CNN-BiLSTM و XGBoost	نظرهای متون فارسی	برت، CNN-BiLSTM، XGBoost	۹۳/۷۴ درصد
لی و همکاران [۴۶] (۲۰۲۴)	برت برای تحلیل لحن و احساسات	نظرهای شبکه‌های اجتماعی در پلتفرم‌های جدید	برت	تشخیص دقیق نگرش‌های احساسی
طلعت [۴۷] (۲۰۲۳)	ترکیب برت با BiLSTM و BiGRU	سه مجموعه داده مختلف	برت، BiLSTM، BiGRU	بهبود نسبت به برت خالص
مانوئل-ایلی و همکاران [۴۸] (۲۰۲۳)	تحلیل جامع با برت	متون با کاربردهای مختلف	برت	دقت بالا در کاربردهای مختلف
پژوهش حاضر	تحلیل احساسات جنبه‌محور تلفن همراه با روش‌های متوازن‌سازی و افزایش داده‌ها	۶۱۸۰ نظر درباره تلفن همراه از دیجی کالا	CNN، برت با تنظیم اختصاصی، فست‌تکست	بهبود دقت نسبت به کارهای گذشته و متوازن‌سازی داده‌ها

۳- روش پیشنهادی

در این بخش روش پیشنهادی برای تحلیل احساسات روی نظرهای کاربران با استفاده از روش‌های یادگیری عمیق و مدل‌های ارائه‌شده برای زبان فارسی بیان شده است. روند کلی این پژوهش، در شکل ۱ نمایش داده شده است.

۳-۱- پیش‌پردازش داده‌ها

یکی از اولین گام‌ها در مسائل پردازش زبان طبیعی، پیش‌پردازش داده‌ها است و برای دستیابی به عملکرد بهتر مدل‌های یادگیری، از اهمیت زیادی برخوردار است. کارهای انجام‌شده برای پیش‌پردازش داده‌ها به شرح زیر است:

متن، اصلاح فاصله‌ها در پیشوندها و پسوندها، حذف اعراب، حذف تکرارهای زائد حروف، جایگزینی برخی از حروف و نشانه‌ها با حروف و نشانه‌های فارسی و جایگزینی اعداد انگلیسی با معادل فارسی انجام می‌گردد.

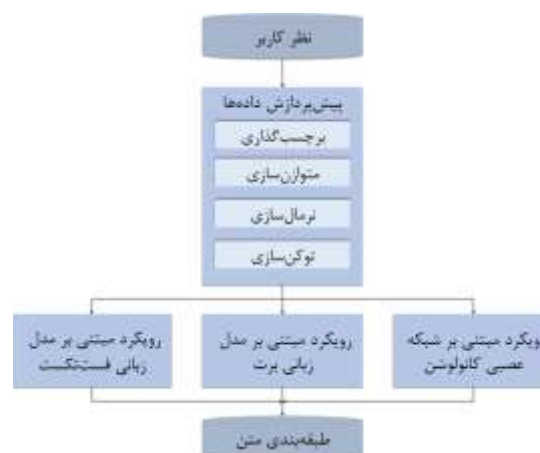
۳-۱-۴- توکن‌سازی

توکن‌سازی مرحله‌ای اساسی در پردازش زبان طبیعی است که متن را به واحدهای معنادار تقسیم می‌کند. این واحدها شامل کلمات، نشانه‌های نگارشی، اعداد و یا واحدهای معنادار دیگر در زبان است. در توکن‌سازی چند نکته حائز اهمیت است، این موارد عبارتند از:

- حفظ معنا: هدف اصلی توکن‌سازی، حفظ معنا در متن است. توکن‌ها باید واحدهای معنادار متن را به نحوی نمایش دهند که تحلیل صحیح از آن‌ها امکان پذیر باشد.
- کار با علائم نگارشی: نشانه‌های نگارشی به‌عنوان توکن‌های جداگانه در نظر گرفته می‌شوند. این امر برای کارهایی مانند تحلیل احساسات، که نشانه‌ها می‌توانند بر تحلیل احساس جمله مؤثر باشند، مهم است.
- پیشوندهای غیر قابل شکستن: در زبان فارسی، پیشوندهایی مثل "می" در "می‌خواهم" غیرقابل شکستن هستند و باید به کلمه‌ای دیگر چسبیده باشند.
- کنترل فضاها: توکن‌سازی باید فضاهای غیرضروری را حذف و متن را در مکان‌های مناسب، بین کلمات و نشانه‌های نگارشی، تقسیم کند.

۳-۲- رویکرد مبتنی بر شبکه عصبی پیچشی

برای تحلیل احساسات با استفاده از شبکه عصبی پیچشی، پس از توکن‌سازی داده‌ها و تبدیل متن به عدد، از لایه جاسازی^{۳۹} برای تبدیل این اعداد به بردارهای عددی استفاده می‌شود. این لایه برای کار با داده‌های متنی استفاده می‌شود. سپس با استفاده از لایه جاسازی، این شناسه‌ها به بردارهای عددی با ابعاد مشخص تبدیل می‌شوند. در مرحله بعد، با اعمال لایه‌های پیچشی با اندازه‌ها و تعداد فیلترهای



شکل ۱- رویکرد پیشنهادی برای تحلیل احساسات مبتنی بر جنبه

۳-۱-۱- برچسب‌گذاری

داده‌ها در این پژوهش، نظرهای مربوط به تلفن همراه است. این نظرها به سه دسته‌بندی مختلف بر اساس جنبه‌های مربوط به دوربین، باتری و قیمت تقسیم و جداگانه برچسب‌گذاری شده است. با توجه به نقاط قوت و ضعف مطرح‌شده توسط کاربر و یک سری کلمات کلیدی مربوط به جنبه‌های انتخاب‌شده، برچسب‌گذاری نظرها در سه کلاس مثبت، منفی و خنثی انجام شده است. برای مثال جمله‌ی "این گوشی قیمت مناسبی دارد و کیفیت باتری خوب است"، در جنبه‌های قیمت و باتری دارای برچسب مثبت و در جنبه دوربین دارای برچسب خنثی است.

۳-۱-۲- متوازن‌سازی

متوازن‌سازی داده‌ها در مدل‌های یادگیری ماشین، هنگامی که داده‌هایی با کلاس‌های نامتوازن داریم، از اهمیت زیادی برخوردار است. ضرورت انجام متوازن‌سازی داده‌ها در مدل‌های یادگیری عمیق برای جلوگیری از بیش‌برازش^{۳۸}، عملکرد بهتر در کلاس‌های با فراوانی کمتر و تعادل در ارزیابی است.

۳-۱-۳- نرمال‌سازی

یکی از مراحل ضروری در پیش‌پردازش داده‌ها، هنگام کار با داده‌های متنی، نرمال‌سازی متن است. این مرحله شامل تبدیل متن به فرمی استاندارد و متعارف برای افزایش دقت در پردازش زبان طبیعی است. در متن فارسی، به علت وجود اشکال گوناگون در نوشتار، برخی از کلمات یا حروف به چندین شکل مختلف نمایش داده می‌شوند. در نرمال‌سازی

³⁹ Embedding Layer

³⁸ Overfitting

فرآیند انطباق مدل برت از قبل آموزش دیده برای یک کار خاص و تنظیم پارامترهای آن با داده‌های جدید است. این فرآیند شامل افزودن یک لایه در بالای رمزگذار و آموزش مدل با یک تابع هزینه و بهینه‌ساز مناسب است. برای تحلیل احساسات با استفاده از مدل برت، پس از پیش‌پردازش داده‌ها، عملیات توکن‌سازی با استفاده از توکن‌ساز برت انجام می‌شود. این توکن‌سازی کمک می‌کند تا متن به توکن‌های واحد تبدیل شود و برای ورودی مدل برت قابل قبول باشد. پس از آن از یک مدل برت از قبل آموزش دیده استفاده می‌شود و مدل برت روی داده‌ها آموزش داده می‌شود. در نهایت مدل آموزش دیده، ارزیابی و معیارهای انتخاب شده روی داده‌های آزمایشی برای ارزیابی عملکرد مدل محاسبه می‌شود.

۳-۴- رویکرد مبتنی بر مدل زبانی فست‌تکست

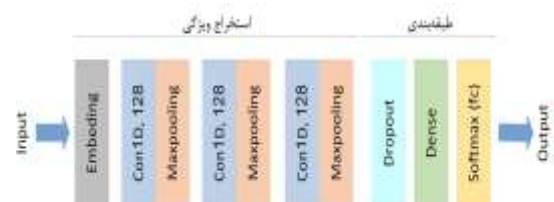
برای مدل فست‌تکست، با استفاده از داده‌های پیش‌پردازش شده، مدل آموزش داده می‌شود. در این مرحله، مدل روی داده‌های آموزشی، آموزش داده می‌شود و وزن-های مربوط برای کلمات استخراج می‌شوند. با بهره‌گیری از الگوریتم Word2Vec، برای ایجاد جاسازی کلمات و با استفاده از معانی مشترک کلمات، بردارهای عددی برای هر کلمه ایجاد می‌شود. پس از ساخت این بردارها، مدل برای طبقه‌بندی متن آموزش داده می‌شود. داده‌های مورد استفاده برای آموزش باید شامل برچسب‌ها و داده‌های آموزش در یک قالب مشخص باشد. برای آموزش مدل پارامترهای مختلف از جمله نرخ یادگیری^{۴۲}، تعداد دوره‌های^{۴۳} آموزش و مقدار n-gram قابل تنظیم است. همچنین از تابع هزینه مناسب برای طبقه‌بندی متن در فرآیند آموزش استفاده می‌شود. پس از آموزش، مدل ارزیابی و میزان دقت و عملکرد آن روی داده‌های آزمایشی محاسبه می‌شود.

۴- پیاده‌سازی و ارزیابی

۴-۱- فراهم‌سازی داده‌ها

داده‌های مورد استفاده در اینجا شامل نظرهای فارسی کاربران در مورد تلفن همراه، به همراه نقاط قوت و ضعف محصول و برچسب "توصیه می‌کنم یا توصیه نمی‌کنم" برای

مختلف، عملیات پیچشی روی بردارهای ورودی انجام شده و ویژگی‌های مختلف استخراج می‌شود. انتخاب چندین لایه پیچشی با فیلترهای افزایشی، برای دستیابی به ویژگی‌های سلسله‌مراتبی در داده‌های ورودی دارای اهمیت است. هر لایه پیچشی دارای اندازه فیلتر ۱۲۸ است و از تابع فعال‌سازی Relu استفاده می‌کند. لایه Conv1D یک نوع لایه پیچشی است که برای پردازش داده‌های یک‌بعدی (مانند پردازش متن و دنباله‌ها) استفاده می‌شود و الگوها و ویژگی‌های مهم در دنباله‌های متنی مانند تشخیص احساسات را استخراج می‌کند. پس از هر لایه پیچشی، عملیات ادغام در خروجی لایه‌های پیچشی برای انتخاب ویژگی‌های مهم از بردارهای پیچشی صورت می‌گیرد. در نهایت از لایه‌های کاملاً متصل برای ترکیب و تبدیل ویژگی‌های استخراج شده به خروجی نهایی استفاده می‌شود. لایه چگال^{۴۰}، یک لایه کاملاً متصل در شبکه عصبی عمیق است. این لایه به‌عنوان لایه پردازش اصلی در شبکه عمل می‌کند و وظیفه تبدیل ورودی‌ها به خروجی‌های پیش‌بینی شده را دارد. در واقع این لایه برای بهبود ویژگی‌ها و کاهش بُعد، به مرحله طبقه‌بندی نهایی اضافه می‌شود و نقش حیاتی در تجمیع دانش به‌دست آمده از لایه‌های قبلی را ایفا می‌کند. در آخرین مرحله، از یک لایه چگال با تابع فعال‌سازی Softmax استفاده شده تا احتمالات خروجی برای طبقه‌بندی نهایی محاسبه شود. همچنین، در روند ساخت مدل پس از مرحله استخراج ویژگی، از یک لایه حذف تصادفی^{۴۱} جهت جلوگیری از بیش‌برازش مدل و افزایش عمومیت آن استفاده شده است. در شکل ۲ ساختار شبکه‌ی عصبی پیچشی پیشنهادی نمایش داده شده است.



شکل ۲- ساختار شبکه عصبی پیچشی پیشنهادی

۳-۳- رویکرد مبتنی بر مدل زبانی برت

مدل برت یک مدل از قبل آموزش دیده است که روی مجموعه بزرگی از داده‌ها آموزش داده می‌شود، سپس برای یک کار خاص تنظیم دقیق می‌شود. تنظیم دقیق برت،

⁴² Learning Rate

⁴³ Epoch

⁴⁰ Dense

⁴¹ Dropout

اولیه (با برچسب دستی) و داده‌های نهایی (ترکیب داده‌ها با برچسب‌های دستی و خودکار، برای ارزیابی رویکردهای پیشنهادی استفاده شده است).

جدول ۲- تعداد برچسب‌ها برای هر جنبه با روش

برچسب‌گذاری دستی

جنبه	مثبت	منفی	خنثی
دوربین	۳۴۰	۲۰۰	۱۲۶۷
باتری	۳۴۲	۱۱۷	۱۳۴۸
قیمت	۱۰۸۳	۱۳۵	۵۸۹

جدول ۳- تعداد برچسب‌ها برای هر جنبه با روش برچسب-

گذاری خودکار

جنبه	مثبت	منفی	خنثی
دوربین	۲۷۷۴	۱۱۸۴	۴۱۵
باتری	۳۱۶۸	۵۷۱	۶۳۴
قیمت	۱۳۲۳	۳۳۸	۲۷۱۲



شکل ۳- روش برچسب‌گذاری داده‌ها به صورت خودکار.

جدول ۴- تعداد برچسب‌های مثبت، منفی و خنثی در هر

دسته‌بندی

دسته‌بندی	مثبت	منفی	خنثی
دوربین	۱۳۱۴	۱۳۸۴	۱۶۸۲
باتری	۳۵۱۰	۶۸۸	۱۹۸۲
قیمت	۲۴۰۶	۴۷۳	۳۳۰۱

۴-۳-۱- نرمال‌سازی

نرمال‌سازی با استفاده از کتابخانه‌ی هضم، برای پردازش زبان فارسی، به تغییرات رایج در خط فارسی مانند اشکال مختلف کاراکترها یا نشانه‌ها می‌پردازد. متن نرمال‌شده حاوی تغییراتی همچون اصلاح فاصله‌گذاری‌ها و نیم‌فاصله‌ها، حذف تکرارهای زائد حروف و تبدیل اعداد لاتین به اعداد فارسی است. برای مثال، جمله‌ی "دوربین این گوشی خیلییییی خوب است و نسبت به قیمت آن می‌ارزد"، پس از نرمال‌سازی به جمله‌ی "دوربین این گوشی خیلی خوب

هر نظر است. داده‌ها از وب سایت دیجی کالا^{۴۴}، به عنوان یک بازارگاه اینترنتی، جمع‌آوری شده است. برای انجام عملیات پیش‌پردازش، از کتابخانه‌ی هضم^{۴۵} استفاده شده است. هضم یک کتابخانه‌ی پایتون برای پردازش زبان طبیعی روی متن فارسی است. از هضم برای نرمال‌سازی متن، توکن‌سازی^{۴۶}، یافتن ریشه کلمات^{۴۷}، تحلیل نحوی و صرفی جملات و شناسایی وابستگی‌های دستوری استفاده می‌شود.

۴-۲- برچسب‌گذاری داده‌ها

برچسب‌گذاری داده‌های مورد استفاده در این پژوهش با دو روش انجام شده است. در روش اول، تعداد ۱۸۰۷ مورد از نظریات کاربری به صورت دستی با توجه به جنبه‌های دوربین، باتری و قیمت در سه کلاس مثبت، منفی و خنثی برچسب‌گذاری شد. در جدول ۲، تعداد برچسب‌ها برای هر جنبه با روش برچسب‌گذاری دستی به تفکیک بیان شده است. به علت زمان‌بر بودن برچسب‌گذاری داده‌ها به صورت دستی و همچنین تعداد کم داده‌های برچسب‌گذاری شده، جهت بهبود عملکرد رویکردهای پیشنهادی تصمیم گرفته شد که از یک روش برچسب‌گذاری خودکار استفاده شود. برچسب‌گذاری خودکار داده‌ها با توجه به نقاط قوت و ضعف مطرح‌شده توسط کاربر و کلمات کلیدی مرتبط با هر جنبه انجام شده است. بنابراین، هر یک از نقاط قوت و ضعف مطرح‌شده که مربوط به جنبه‌های مورد نظر است، با یک کلمه با بار احساسی مثبت یا منفی، جهت تکمیل مفهوم جمله، به متن اصلی نظر اضافه کرده و با توجه به آن برچسب‌گذاری شده است. با استفاده از این روش، تعداد ۴۳۷۳ نظر برچسب‌گذاری شده است. در شکل ۳ نمونه‌ای از روند برچسب‌گذاری خودکار داده‌ها نشان داده شده است. همچنین در جدول ۳ تعداد برچسب‌ها برای هر جنبه با روش برچسب‌گذاری خودکار به تفکیک بیان شده است.

۴-۳-۲- مشخصات داده‌ها

پس از فرآیند برچسب‌گذاری، داده‌های نهایی در مجموع شامل ۶۱۸۰ نظر در مورد تلفن همراه است. داده‌ها در دسته‌بندی‌های دوربین، قیمت و باتری برچسب‌گذاری شده‌اند و تعداد برچسب‌های مربوط به هر دسته‌بندی در کلاس‌های مثبت، منفی و خنثی در جدول ۴ مطرح شده است. در این پژوهش از دو مجموعه داده، شامل داده‌های

⁴⁶ Tokenization

⁴⁷ Lemmatization

⁴⁴ <https://www.digikala.com>

⁴⁵ Hazm

دوربین	۳۱۱۴	۲۷۰۱	۳۳۴۸	۶۱۸۰	۹۱۶۳
باتری	۳۵۱۰	۳۲۸۰	۳۹۴۰	۶۱۸۰	۱۰۷۳۰
قیمت	۲۴۰۶	۳۲۰۳	۳۳۰۱	۶۱۸۰	۸۹۱۰

۴-۴- مشخصات سیستم

برای پیاده سازی رویکردهای پیشنهادی در این پژوهش از سرویس گوگل کولب استفاده شده است. در این پژوهش، سرویس مورد استفاده برای پیاده سازی دارای ۱۲ گیگابایت حافظه رم، ۱۰۰ گیگابایت فضای ذخیره سازی، دو پردازنده Intel Xeon و پردازنده گرافیکی NVIDIA Tesla است.

۴-۵- تحلیل و ارزیابی رویکردهای پیشنهادی

در این پژوهش برای پیاده سازی رویکردهای پیشنهادی، در تمام اجراها از تابع هزینه Categorical Cross Entropy مناسب برای مسائل چندکلاسه، بهینه ساز Adam و نرخ یادگیری ۰/۰۰۰۰۱ استفاده شده است. همچنین نسبت داده های آزمون و داده های ارزیابی ۰/۳ و مقدار k-fold، ۵ در نظر گرفته شده است. تعداد دوره ها نیز مقادیر ۱، ۵ و ۱۰ لحاظ شده است. در ادامه، معیارهای ارزیابی مورد استفاده در این پژوهش، جزییات مربوط به پیاده سازی هر کدام از رویکردها و نتایج حاصل از آنها با داده های اولیه و نهایی بیان شده است.

۴-۵-۱- معیارهای ارزیابی

در این پژوهش برای ارزیابی عملکرد رویکردهای پیشنهادی از دو معیار ارزیابی صحت^{۵۲} و امتیاز F1^{۵۳} استفاده شده است. در روابط ۱ تا ۴ نحوه محاسبه این معیارها بیان شده است. در این روابط، TP تعداد نمونه هایی است که به درستی در کلاس مثبت و FP تعداد نمونه هایی است که اشتباه در کلاس مثبت تشخیص داده شده اند. همچنین TN تعداد نمونه هایی است که به درستی در کلاس منفی و FN تعداد نمونه هایی است که اشتباه در کلاس منفی تشخیص داده شده اند.

صحت: نشان دهنده درصد نمونه های صحیحی است که مدل تشخیص داده است. این معیار به عنوان نسبت تعداد پیش بینی های صحیح به تعداد کل پیش بینی ها ارائه می شود و به این معناست که مدل تا چه اندازه پیش بینی درست داشته است.

است و نسبت به قیمت آن می ارزد" تبدیل می شود.

۴-۳-۲- توکن سازی

کتابخانه ای هضم دارای تابعی به نام word tokenize است که به طور خاص، برای توکن سازی متن فارسی طراحی شده است. به این تابع متن ورودی داده می شود و خروجی آن لیستی از توکن های تشکیل دهنده متن است. به این ترتیب، از توکن های حاصل برای تحلیل و پردازش متن می توان استفاده کرد.

۴-۳-۳- افزایش داده ها

یکی از چالش ها در رویکردهای پردازش زبان طبیعی، کمبود داده های فارسی است. افزایش داده ها راهی برای مقابله با این چالش است. تصویر را می توان به راحتی با چرخاندن، استفاده از فیلترها و افزودن نویز افزایش داد؛ اما در زمینه کار با داده های متنی، با توجه به پیچیدگی زبان، افزایش داده ها دشوار است. در اینجا از روش های متوازن سازی و افزایش داده ها استفاده شده است. روش افزایش آسان داده ها^{۴۸} از چهار عملیات ساده ولی قدرتمند شامل جایگزینی مترادف^{۴۹}، درج تصادفی^{۵۰}، جابه جایی تصادفی^{۵۱} و حذف تصادفی^{۵۲} تشکیل شده است. برای استفاده از روش افزایش آسان داده ها از کتابخانه nlpaug بهره گرفته شده است. در جدول ۵ نتایج استفاده از هر کدام از این عملیات روی مدل زبانی برت بیان شده است. با توجه به نتایج به دست آمده، از روش جابه جایی تصادفی برای بقیه رویکردها استفاده شده است. در جدول ۶ تعداد داده ها پس از افزایش، با استفاده از عملیات جابه جایی تصادفی، روی داده های نهایی به تفکیک بیان شده است.

جدول ۵- نتایج افزایش داده ها روی مدل زبانی برت

عملیات	دوربین	باتری	قیمت
جایگزینی مترادف	۷۰/۵	۷۷/۴۳	۸۰/۱۰
درج تصادفی	۷۵/۲۰	۸۷/۷۷	۸۲/۶۱
جابه جایی تصادفی	۸۳/۱۲	۹۳/۸۸	۹۰/۳۲
حذف تصادفی	۷۲/۳۰	۸۳/۴۰	۷۹/۹۰

جدول ۶- تعداد برچسب ها پس از افزایش داده ها با روش

جابه جایی تصادفی

جنبه	مثبت	منفی	خنثی	اولیه	تعداد نهایی پس از افزایش
------	------	------	------	-------	--------------------------

⁵² Random Deletion

⁵³ Accuracy

⁵⁴ F1-Score

⁴⁸ Easy Data Augmentation (EDA)

⁴⁹ Synonym

⁵⁰ Random Insertion

⁵¹ Random Swap

دوربین	بله	۹۷/۸۳	۸۱/۷۹
باتری	خیر	۸۵/۶۷	۶۸/۱۵
	بله	۹۷/۴۶	۸۲/۷۵
قیمت	خیر	۹۳/۲۰	۶۶/۲۳
	بله	۹۸	۸۱/۱۰

در جدول ۹ نتایج رویکرد شبکه عصبی پیچشی روی داده‌های نهایی با ۵ و ۱۰ دوره گزارش شده است. نتایج حاصل نشان می‌دهد که در مجموع رویکرد شبکه عصبی پیچشی، روی داده‌های افزایش‌یافته با ۱۰ دوره بهترین عملکرد را در هر سه جنبه داشته است. بنابراین، می‌توان نتیجه گرفت که روش افزایش آسان داده‌ها با عملیات جابه‌جایی تصادفی، برای متوازن‌سازی داده‌ها، تأثیر مثبتی در عملکرد مدل داشته و باعث بهبود نتایج شده است.

جدول ۹- نتایج شبکه عصبی پیچشی روی داده‌های نهایی

دوره		۵		۱۰	
جنبه	متوازن	صح	امتیاز F	صحت	امتیاز F
دوربین	خیر	۷/۵۸	۶۹/۹۲	۹۵/۹۰	۷۵/۴۲
	بله	۹/۸۲	۷۷/۴۱	۹۷/۹۲	۷۷/۴۵
باتری	خیر	۸/۴۵	۶۱/۴۵	۸۷/۰۹	۷۱/۳۱
	بله	۹/۸۵	۷۹/۹۹	۹۷/۴۳	۸۰/۱۷
قیمت	خیر	۸/۴۸	۸۴/۷۴	۹۵/۵۱	۸۰/۸۹
	بله	۹/۲۵	۷۸/۴۸	۹/۶۶	۸۸/۹۴

در شکل‌های ۴، ۵ و ۶ نمودار ارزیابی دقت و هزینه شبکه عصبی پیچشی برای هر سه جنبه روی داده‌های نهایی و متوازن‌شده نمایش داده شده است. همانطور که در هر سه نمودار ارزیابی دقت مشخص است، در طی دوره‌های ابتدایی، اختلاف بین دقت اعتبارسنجی و دقت آموزش نشان می‌دهد که بیش‌برازش رخ داده است. پس از سومین دوره، دقت اعتبارسنجی شروع به همگرایی با دقت آموزش می‌کند. این همگرایی نشان‌دهنده تثبیت آموزش مدل است و مدل به‌عنوان یک تعمیم‌دهنده به خوبی عمل کرده است.

فراخوانی^{۵۵}: نسبت تعداد نمونه‌های درست تشخیص داده‌شده به تعداد کل نمونه‌های واقعی در یک کلاس را نشان می‌دهد.

دقت^{۵۶}: نسبت تعداد نمونه‌های واقعی که به درستی در یک کلاس مشخص تشخیص داده شده‌اند.

$$Accuracy = \frac{TP + TN}{TP} \quad (۱)$$

$$Recall = \frac{TP}{TP + FN} \quad (۲)$$

$$Precision = \frac{TP}{TP + FP} \quad (۳)$$

$$F1-Score = 2 * \frac{Precision * Recall}{Precision + Recall} \quad (۴)$$

۴-۵-۲- پیاده‌سازی و ارزیابی رویکرد شبکه عصبی پیچشی

برای پیاده‌سازی شبکه عصبی پیچشی برای طبقه‌بندی متن، از کتابخانه Keras و Tensorflow استفاده شده است. با بهره‌گیری از توابع فعال‌سازی Relu و Softmax و تنظیم ابرپارامترهای مربوط، مدل آموزش داده و ارزیابی شده است. در جدول ۷ مقادیر ابرپارامترهای مربوط به شبکه عصبی پیچشی آمده است.

جدول ۷- ابرپارامترهای شبکه عصبی پیچشی

پارامتر	مقدارهای آزمایش شده	مقدار بهینه انتخاب شده
سایز فیلتر	۳، ۴، ۵	۳
اندازه فیلتر	۳۲، ۶۴، ۱۲۸	۱۲۸
اندازه دسته	۶۴، ۱۲۸، ۲۵۶	۶۴
نرخ حذف تصادفی	۰/۳، ۰/۵، ۰/۷	۰/۵

در جدول ۸ نتایج رویکرد شبکه عصبی پیچشی روی داده‌های اولیه با ۱ دوره گزارش شده است. منظور از متوازن همان روش افزایش داده‌ها است که در صورت اعمال روی داده‌ها با "بله" و در صورت عدم اعمال روی داده‌ها با "خیر" مشخص شده است. نتایج حاصل نشان می‌دهد که به‌صورت کلی، رویکرد مورد نظر روی داده‌های افزایش‌یافته بهترین عملکرد را در هر سه جنبه داشته است.

جدول ۸- نتایج شبکه عصبی پیچشی روی داده‌های اولیه

جنبه	متوازن	صحت	امتیاز F1
	خیر	۸۰/۷۱	۶۱/۶۶

زبان فارسی و مبتنی بر ترنسفورمر، داده‌های اولیه و نهایی روی مدل ارزیابی شده است. در جدول ۱۱ نتایج رویکرد برت روی داده‌های اولیه با ۵ دوره گزارش شده است. نتایج حاصل نشان می‌دهد که رویکرد برت، روی داده‌های افزایش یافته بهترین عملکرد را در هر سه جنبه داشته است. در جدول ۱۲ نتایج رویکرد برت روی داده‌های نهایی با ۱ و ۵ دوره گزارش شده است. نتایج حاصل از این رویکرد نیز نشان می‌دهد که در مجموع رویکرد برت، روی داده‌های افزایش یافته با ۵ دوره بهترین عملکرد را در هر سه جنبه داشته است. بنابراین، می‌توان نتیجه گرفت که روش افزایش آسان داده‌ها با عملیات جابه‌جایی تصادفی، برای متوازن سازی داده‌ها، تأثیر مثبتی روی عملکرد این مدل نیز داشته و باعث بهبود نتایج شده است.

جدول ۱۰- ابرپارامترهای مدل برت

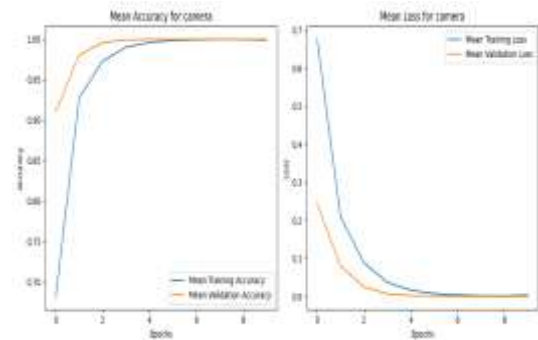
مقدار بهینه انتخاب شده	مقدارهای آزمایش شده	ابرپارامتر
۶۴	۶۴، ۱۲۸، ۲۵۶	اندازه دسته
۲۵۶	۱۲۸، ۲۵۶، ۵۱۲	حداکثر طول ورودی ^{۵۸}
۰/۵	۰/۳، ۰/۵، ۰/۷	نرخ حذف تصادفی

جدول ۱۱- نتایج مدل زبانی برت روی داده‌های اولیه

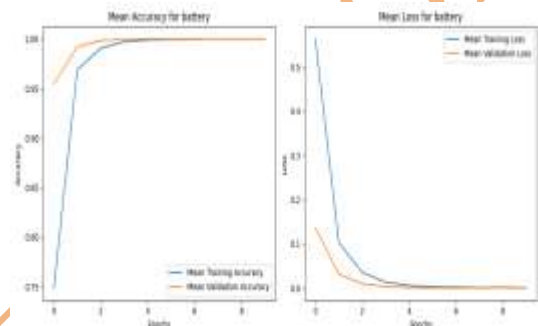
جنبه	متوازن	صحت	امتیاز F1
دوربین	خیر	۷۳/۴۸	۶۸/۷۸
	بله	۸۳/۱۲	۸۲/۹۶
باتری	خیر	۷۶/۱۱	۶۸/۴۵
	بله	۹۳/۸۸	۹۳/۱۰
قیمت	خیر	۷۷/۲۷	۷۴/۸۳
	بله	۹۰/۳۲	۹۰/۳۱

جدول ۱۲- نتایج مدل زبانی برت روی داده‌های نهایی

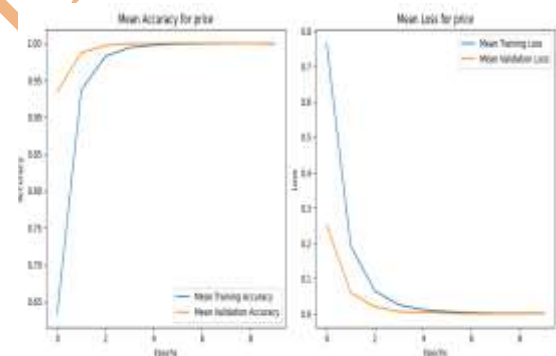
دوره		۱		۵	
جنبه	متوازن	صحت	امتیاز F1	صحت	امتیاز F1
دوربین	خیر	۸/۲۹	۸۸/۸۸	۸/۷۶	۸۹/۶۸
	بله	۹/۳۴	۹۰/۳۷	۹/۸۶	۹۳/۰۴
باتری	خیر	۹/۹۳	۹۱/۷۱	۹/۴۳	۹۲/۱۴
	بله	۹/۹۳	۹۱/۷۱	۹/۴۳	۹۲/۱۴



شکل ۴- نمودار ارزیابی دقت و هزینه شبکه عصبی پیچشی در دسته‌بندی دوربین روی داده‌های نهایی



شکل ۵- نمودار ارزیابی دقت و هزینه شبکه عصبی پیچشی در دسته‌بندی باتری روی داده‌های نهایی



شکل ۶- نمودار ارزیابی دقت و هزینه شبکه عصبی پیچشی در دسته‌بندی قیمت روی داده‌های نهایی

۴-۵-۳- پیاده‌سازی و ارزیابی رویکرد برت

در رویکرد برت از مدلی از پیش آموزش دیده روی نظرهای کاربران دیجی کالا، در بستر هاگینگ فیس^{۵۷} استفاده شده است. هاگینگ فیس بستری آموزشی برای پژوهش در زمینه پردازش زبان طبیعی است. برای استفاده از مدل برت از کتابخانه Transformers استفاده شده است. در جدول ۱۰ مقادیر ابرپارامترهای مربوط به مدل برت آمده است. ادامه با استفاده از یک مدل برت از پیش آموزش دیده برای

⁵⁸ Max Sequence Length⁵⁷ Hugging Face

با بررسی روش‌هایی برای متوازن‌سازی داده‌ها، تأثیر آن در نتایج هر رویکرد بررسی شد. مدل زبانی برت به‌عنوان یک معماری قدرتمند در پردازش زبان طبیعی شناخته شده است و در بسیاری از رویکردهای حوزه زبان طبیعی کارآمد است. به دلیل پیچیدگی معماری برت و تعداد بالای پارامترها، سرعت اجرای آن کمتر از روش‌های فست‌تکست و شبکه عصبی پیچشی است. بنابراین، آموزش و استفاده از مدل برت نیازمند زمان و منابع محاسباتی بیشتری است. در این پژوهش، با استفاده از رویکرد برت، بهترین نتایج را از نظر امتیاز F1 در هر سه جنبه داشت؛ با این حال زمان اجرا برای آموزش و ارزیابی در هر جنبه بیش از ۱۶۰ دقیقه شده است. شبکه عصبی پیچشی نیز به‌عنوان روشی قدرتمند در حوزه پردازش زبان طبیعی و کارهایی مانند تحلیل احساسات و دسته‌بندی متن مورد استفاده قرار می‌گیرد. معمولاً این رویکرد، به دلیل ساختار ساده‌تر و تعداد کمتر پارامترها، سرعت اجرای بالاتری نسبت به رویکرد برت دارد. بنابراین، شبکه عصبی پیچشی در مقیاس بزرگ‌تر و با داده‌های بیشتر، می‌تواند سریع‌تر اجرا شود. در این پژوهش، با استفاده از این رویکرد، از نظر معیار صحت، به بهترین نتایج در هر سه جنبه دست یافتیم. زمان اجرا برای آموزش و ارزیابی برای هر جنبه نیز در بازه زمانی ۱۰۰ الی ۱۲۰ دقیقه قرار گرفته که نسبت به مدل برت زمان مطلوب‌تری برای رسیدن به خروجی است. مدل زبانی فست‌تکست روشی سبک برای پردازش زبان طبیعی است و برای کارهایی مانند بردارسازی کلمات، تحلیل احساسات و طبقه‌بندی متن استفاده می‌شود. این روش با داده‌های با حجم زیاد به خوبی کار می‌کند و نتایج قابل قبولی را با سرعت بالا ارائه می‌دهد. در این پژوهش، با استفاده از این رویکرد نیز نتایج خوبی در هر سه جنبه به‌دست آمد. مزیت استفاده از این روش، زمان اجرای بسیار کم آن است و زمان اجرا برای آموزش و ارزیابی برای هر جنبه کمتر از ۱۰ ثانیه شده است که نسبت به دو رویکرد قبلی، زمان بسیار خوبی برای رسیدن به خروجی است. مدل‌های شبکه عصبی پیچشی، برت و فست‌تکست در بهترین حالت به ترتیب مقادیر ۹۸/۹۴، ۹۶/۳۶ و ۸۹/۳۰ درصد را با معیار امتیاز F1 کسب کردند.

۵- نتیجه‌گیری و پیشنهادها

مردم در سراسر جهان، نظرات خود را در مورد موضوعات

۹۶/۳۶	۹۵/۱۱	۹۳/۷۱	۹/۵۲ ۳	بله	
۸۹/۲۰	۹/۱۳ ۰	۸۹/۱۸	۸/۸۵ ۹	خیر	قیمت
۹۳/۴۶	۹/۷۷ ۲	۹۰/۸۶	۹/۵۵ ۰	بله	

۴-۵-۴- پیاده‌سازی و ارزیابی رویکرد فست‌تکست

برای پیاده‌سازی مدل فست‌تکست از کتابخانه FastText استفاده شده است. با تنظیم ابرپارامترهای مربوط، مدل روی داده‌ها آموزش داده و ارزیابی شده است. این فرآیند شامل ارتقاء وزن‌های مدل و تنظیم آن‌ها به گونه‌ای است که مدل، متن را به درستی دسته‌بندی کند. برای این منظور مقدار n-gram ۲ و ۳ در نظر گرفته شده که مقدار ۳ نتایج بهتری به همراه داشته است. در جدول ۱۳ نتایج رویکرد فست‌تکست روی داده‌های نهایی با ۱ و ۵ دوره گزارش شده است. بر اساس نتایج، رویکرد فست‌تکست نیز همانند دو رویکرد دیگر، روی داده‌های افزایش‌یافته بهترین عملکرد را در هر سه جنبه داشته است. بنابراین، می‌توان نتیجه گرفت که روش افزایش‌آسان داده‌ها تأثیر مثبتی روی عملکرد این مدل نیز داشته و باعث بهبود نتایج شده است.

جدول ۱۳- نتایج مدل زبانی فست‌تکست روی داده‌های

نهایی

دوره		۱		۵	
جنبه	متوازن	صح ت	امتیاز F	صح ت	امتیاز F
دوربین	خیر	۶/۳۴ ۹	۶۰/۷۳	۷/۷۸ ۲	۶۷/۷۰
	بله	۷/۶۱ ۸	۷۲/۱۷	۸/۶۷ ۲	۸۰/۴۲
باتری	خیر	۷/۲۳ ۰	۶۹/۵۰	۷/۶۳ ۵	۷۰/۷۷
	بله	۸/۶۳ ۲	۷۷/۹۷	۸/۷۹ ۵	۸۴/۳۰
قیمت	خیر	۷/۲۷ ۴	۷۲/۵۲	۷/۵۴ ۵	۷۳/۰۸
	بله	۸/۸۳ ۳	۷۶/۱۵	۹/۱۳ ۰	۸۹/۳۰

۴-۶- جمع‌بندی و مقایسه رویکردها

در این پژوهش، از سه رویکرد مختلف تحلیل احساسات و طبقه‌بندی متن از نظر احساسات مثبت، منفی و خنثی در سه جنبه دوربین، باتری و قیمت استفاده شد. همچنین

هر جنبه به دست آمد. این نتایج حاکی از آن است که مدل‌های مورد نظر در پیش‌بینی احساسات نظرهای به-خوبی عمل کرده‌اند. همچنین نتایج نشان می‌دهد که شبکه عصبی پیچشی، از نظر معیار صحت، عملکرد بهتری نسبت به رویکردهای دیگر داشته است. از نتایج به دست آمده در این پژوهش می‌توان در حوزه‌های مختلف از جمله بازاریابی، خدمات مشتریان و توسعه محصولات بهره برد.

۶- کار آینده

به دلیل محدودیت در دسترسی به مجموعه داده‌های متنوع، در آینده با فراهم شدن مجموعه داده‌های بیشتر، مقایسه مدل‌ها روی دامنه‌های متنوع‌تر انجام خواهد شد. همچنین می‌توان با افزودن لایه‌های عمیق‌تر به شبکه‌های پیچشی به دقت بیشتری در تشخیص احساسات دست پیدا کرد. همچنین می‌توان با استفاده از داده‌های بیشتر و در نظر گرفتن جنبه‌های بیشتر، به استخراج برخی از جنبه‌های کیفی و تمرکز روی یک سری کالاهای خاص پرداخت و تأثیر پارامترهایی همچون گذر زمان و نوع کالا را در رویکردهای تحلیل احساسات بررسی کرد. با تحلیل احساسات نظرها و بازخوردهای کاربران و درک بهتر از نیازها و خواسته‌های آن‌ها، می‌توان به صورت شخصی سازی-شده پیشنهادهایی را برای آن‌ها ارائه داد. به عنوان مثال، با توجه به نتایج تحلیل احساسات مشتریان، می‌توان محصولات مرتبط را برای جلب رضایت آن‌ها پیشنهاد کرد.

تأییدیه اخلاقی

نویسندگان متعهد می‌شوند که مطالب این مقاله را در هیچ مجله دیگری به چاپ نرسانده‌اند.

تعارض منافع

نویسندگان اعلام می‌کنند که در مورد انتشار این مقاله تعارض منافع وجود ندارد.

مختلفی همچون فیلم‌ها، محصولات، سیاست و مسائل روز دنیا در بسترهای اینترنتی و رسانه‌های اجتماعی به اشتراک می‌گذارند. تحلیل احساسات برای طبقه‌بندی خودکار این احساسات به عنوان احساس مثبت، منفی و یا خنثی استفاده می‌شود. در زمینه تحلیل احساسات تاکنون تحقیقات زیادی روی داده‌های انگلیسی انجام شده؛ اما برای زبان فارسی کارهای کمتری صورت گرفته است. در این پژوهش تلاش شد تا با به کارگیری روش‌های مناسب برای پیش پردازش داده‌های فارسی و استفاده از رویکردهای شبکه عصبی پیچشی و مدل‌های زبانی برت و فست‌تکست، تحلیل احساسات نظرهای کاربران در مورد تلفن همراه در جنبه‌های دوربین، باتری و قیمت در وبسایت دیجی کالا انجام شود. تحلیل احساسات با توجه به جنبه‌های مختلف به ما کمک می‌کند تا تحلیل دقیق‌تری از نظرهای کاربران و همچنین درک بهتری از نیازهای مشتریان داشته باشیم. علاوه بر این، منجر به شناخت جنبه‌های گوناگون از یک تلفن همراه مانند کیفیت دوربین، عمر باتری و قیمت آن می‌شود. در همین راستا، داده‌ها مبتنی بر جنبه و بر اساس سه احساس مثبت، منفی و خنثی برچسب گذاری شد. این پژوهش برای درک بهتر عقاید و نظرهای کاربران و همچنین استفاده از این داده‌ها برای پیش‌بینی نظرهای مثبت، منفی و خنثی انجام شده است. برای ارزیابی عملکرد رویکردهای پیشنهادی، از دو مجموعه داده شامل داده‌های اولیه با برچسب گذاری دستی و داده‌های نهایی با ترکیب داده‌ها با برچسب گذاری دستی و خودکار استفاده شده است. یکی از چالش‌های موجود در این پژوهش، عدم توازن بین داده‌ها در کلاس‌های مثبت، منفی و خنثی بوده است. به همین منظور، برای متوازن سازی جنبه‌ها در هر دسته‌بندی، از روش‌های افزایش داده نیز روی داده‌های نهایی بهره گرفته شده است. در نهایت با استفاده از مدل‌های پیشنهادی، دقت بسیار بالایی در تشخیص نظرهای مثبت، منفی و خنثی در

مراجع

- [1] L. Yue, W. Chen, X. Li, W. Zuo, and M. Yin. "A survey of sentiment analysis in social media." *Knowledge and Information Systems* 60 (2019): 617-663.
- [2] Z. Rajabi, and M. R. Valavi. "A survey on sentiment analysis in Persian: a comprehensive system perspective covering challenges and advances in resources and methods." *Cognitive Computation* 13, no. 4 (2021): 882-902.
- [3] R. Asgarneshad, and S.A. Monadjemi. "Persian sentiment analysis: feature engineering, datasets, and challenges." *Journal of applied intelligent systems & information sciences* 2, no. 2 (2021): 1-21.
- [4] M. Shamsfard. "Challenges and opportunities in processing low resource languages: A study on persian." In *International Conference Language Technologies for All (LT4All)*. 2019.

- [5] M. Birjali, M. Kasri, and A. Beni-Hssane. "A comprehensive survey on sentiment analysis: Approaches, challenges and trends." *Knowledge-Based Systems* 226 (2021): 107134.
- [6] J. K. Rout, et al. "A model for sentiment and emotion analysis of unstructured social media text." *Electronic Commerce Research* 18 (2018): 181-199.
- [7] S. Behdenna, F. Barigou, and G. Belalem. "Sentiment analysis at document level." In *Smart Trends in Information Technology and Computer Communications: First International Conference, SmartCom 2016, Jaipur, India, August 6–7, 2016, Revised Selected Papers 1*, p. 159-168. Springer Singapore, 2016.
- [8] A. Yadav, Ashima, and D. K. Vishwakarma. "Sentiment analysis using deep learning architectures: a review." *Artificial Intelligence Review* 53, no. 6 (2020): 4335-4385.
- [9] T. Chen, R. Xu, Y. He, and X. Wang. "Improving sentiment analysis via sentence type classification using BiLSTM-CRF and CNN." *Expert Systems with Applications* 72 (2017): 221-230.
- [10] HH. Do, P.W. Prasad, A. Maag, and A. Alsadoon. "Deep learning for aspect-based sentiment analysis: a comparative review." *Expert systems with applications* 118 (2019): 272-299.
- [11] M. Wankhade, ACS. Rao, and C. Kulkarni. "A survey on sentiment analysis methods, applications, and challenges." *Artificial Intelligence Review* 55, no. 7 (2022): 5731-5780.
- [12] M. Taboada, J. Brooke, M. Tofiloski, K. Voll, and M. Stede. "Lexicon-based methods for sentiment analysis." *Computational linguistics* 37, no. 2 (2011): 267-307.
- [13] V. Bonta, N. Kumares, and N. Janardhan. "A comprehensive study on lexicon based approaches for sentiment analysis." *Asian Journal of Computer Science and Technology* 8, no. S2 (2019): 1-6.
- [14] Q. Ye, Qiang, Z. Zhang, and R. Law. "Sentiment classification of online reviews to travel destinations by supervised machine learning approaches." *Expert systems with applications* 36, no. 3 (2009): 6527-6535.
- [15] G. Paltoglou, and M. Thelwall. "Twitter, myspace, digg: Unsupervised sentiment analysis in social media." *ACM Transactions on Intelligent Systems and Technology (TIST)* 3, no. 4 (2012): 1-19.
- [16] I. Gupta, and N. Joshi. "Enhanced twitter sentiment analysis using hybrid approach and by accounting local contextual semantic." *Journal of intelligent systems* 29, no. 1 (2019): 1611-1625.
- [17] K. Dashtipour, et al. "Exploiting deep learning for Persian sentiment analysis." In *Advances in Brain Inspired Cognitive Systems: 9th International Conference, Xi'an, China*, p. 597-604. Springer International Publishing, 2018.
- [18] L. Zhang, S. Wang, and B. Liu. "Deep learning for sentiment analysis: A survey." *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery* 8, no. 4 (2018): e1253.
- [19] X. Ouyang, P. Zhou, CH. Li, and L. Liu. "Sentiment analysis using convolutional neural network." In *2015 IEEE international conference on computer and information technology; ubiquitous computing and communications; dependable, autonomic and secure computing; pervasive intelligence and computing*, p. 2359-2364. IEEE, 2015.
- [20] H. Gu, Y. Wang, S. Hong, and G. Gui. "Blind channel identification aided generalized automatic modulation recognition based on deep learning." *IEEE Access* 7 (2019): 110722-110729.
- [21] J. Hassannataj Joloudari, et al. "BERT-deep CNN: State of the art for sentiment analysis of COVID-19 tweets." *Social Network Analysis and Mining* 13, no. 1 (2023): 99.
- [22] A. Radford, K. Narasimhan, T. Salimans, and I. Sutskever. "Improving language understanding by generative pre-training." (2018).
- [23] M. Farahani, M. Gharachorloo, M. Farahani, and M. Manthouri. "Parsbert: Transformer-based model for persian language understanding." *Neural Processing Letters* 53 (2021): 3831-3847.
- [24] E. Grave, P. Bojanowski, P. Gupta, A. Joulin, and T. Mikolov. "Learning word vectors for 157 languages." *arXiv preprint arXiv:1802.06893* (2018).
- [25] J. Devlin, MW. 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*, volume 1 (long and short papers), p. 4171-4186. 2019.

- [26] Y. Cao, Z. Sun, L. Li, and W. Mo. "A study of sentiment analysis algorithms for agricultural product reviews based on improved BERT model." *Symmetry* 14, no. 8 (2022): 1604.
- [27] R. Catelli, S. Pelosi, and M. Esposito. "Lexicon-based vs. Bert-based sentiment analysis: A comparative study in Italian." *Electronics* 11, no. 3 (2022): 374.
- [28] A. Joulin, E. Grave, P. Bojanowski, and T. Mikolov. "Bag of tricks for efficient text classification." *arXiv preprint arXiv:1607.01759* (2016).
- [29] A. Amalia, OS Sitompul, EB Nababan, and T. Mantoro. "An efficient text classification using fasttext for bahasa indonesia documents classification." In *2020 International Conference on Data Science, Artificial Intelligence, and Business Analytics*, p. 69-75. IEEE, 2020.
- [30] MR. Hossain, MH. MM. Hoque, and IH. Sarker. "Text classification using convolution neural networks with fasttext embedding." In *International Conference on Hybrid Intelligent Systems*, p. 103-113. Cham: Springer International Publishing, 2020.
- [31] T. Szandała. "Review and comparison of commonly used activation functions for deep neural networks." *Bio-inspired neurocomputing* (2021): 203-224.
- [32] A. Jurek, MD. Mulvenna, and Y. Bi. "Improved lexicon-based sentiment analysis for social media analytics." *Security Informatics* 4 (2015): 1-13.
- [33] Z. Singla, S. Randhawa, and S. Jain. "Sentiment analysis of customer product reviews using machine learning." In *2017 international conference on intelligent computing and control*, p. 1-5. IEEE, 2017.
- [34] F. Ebrahimi Rashed, and N. Abdolvand. "A supervised method for constructing sentiment lexicon in persian language." *Journal of Computer & Robotics* 10, no. 1 (2017).
- [35] Y. Al Amrani, M. Lazaar, and K. E. El Kadiri. "Random forest and support vector machine based hybrid approach to sentiment analysis." *Procedia Computer Science* 127 (2018): 511-520.
- [36] H. Kim, and YS. Jeong. "Sentiment Classification Using Convolutional Neural Networks" *Applied Sciences* 9, no. 11 (2019): 2347. <https://doi.org/10.3390/app9112347>
- [37] S. Soni, SS. Chouhan, and SS. Rathore. "TextConvoNet: A convolutional neural network based architecture for text classification." *Applied Intelligence* 53, no. 11 (2023): 14249-14268.
- [38] B. Jang, M. Kim, G. Harerimana, S. Kang, and JW. Kim. "Bi-LSTM model to increase accuracy in text classification: Combining Word2vec CNN and attention mechanism." *Applied Sciences* 10, no. 17 (2020): 5841.
- [39] W. Zhang, X. Li, Y. Deng, L. Bing, and W. Lam. "A survey on aspect-based sentiment analysis: Tasks, methods, and challenges." *IEEE Transactions on Knowledge and Data Engineering* 35, no. 11 (2022): 11019-11038.
- [40] A. Omar, N. Siyam, A. Abdel Monem, and K. Shaalan. "Aspect-based sentiment analysis using smart government review data." *Applied Computing and Informatics* 20, no. 1/2 (2024): 142-161.
- [41] N. Majumder, R. Bhardwaj, S. Poria et al. "Improving aspect-level sentiment analysis with aspect extraction." *Neural Computing and Applications* (2022): 1-11.
- [42] D. Tang, B. Qin, X. Feng, and T. Liu. "Effective LSTMs for target-dependent sentiment classification." *arXiv preprint arXiv:1512.01100* (2015).
- [43] M. Roustaei, and M. A. Javadzadeh. "Sentiment Analysis of Persian Language Users Based on the Combination of Semantic Relations and Random Forest." In *Proceedings of the 8th National Conference on Modern Studies and Researches in Computer Science, Electrical, and Mechanical Engineering of Iran, Tehran, 2022*. (in Persian)
- [44] M. Ghasempour Marzbali, "Aspect-Based Sentiment Analysis Using BERT." In *Proceedings of the 4th International Conference on Artificial Intelligence and Its Future Perspective in Electrical, Computer, Mechanical, and Telecommunication Engineering, Mashhad, 2024*. (in Persian)
- [45] F. Foroutan, and M. Rabiei. "Applying Deep Learning to Improve Sentiment Analysis Results of Persian Reviews in Online Retail Stores." *Iranian Journal of Information and Communication Technology* 57, no. 15 (2023): 122-137. (in Persian)
- [46] B. Li, X. Liu, and R. Zhang. "Employing the BERT model for sentiment analysis of online commentary." *Applied and Computational Engineering* 32 (2024): 241-247.

[47] A. S. Talaat. "Sentiment analysis classification system using hybrid BERT models." *Journal of Big Data* 10, no. 1 (2023): 110.

[48] D. Manuel-Ilie, P. A. Gabriel, and R. G. Crețulescu. "Sentiment Analysis Using Bert Model." *International Journal of Advanced Statistics and IT&C for Economics and Life Sciences* 13, no. 1 (2023): 59-66.

UNCORRECTED PROOF