

ارائه روشی مبتنی بر یادگیری عمیق و واژه‌نامه حسی برای تحلیل احساسات متون فارسی

سمیرا نوفرستی^{۱*}، دانشیار، مهشید میری^۲، فارغ التحصیل مقطع کارشناسی کامپیوتر

^۱ گروه مهندسی فناوری اطلاعات - دانشکده مهندسی برق و کامپیوتر - دانشگاه سیستان و بلوچستان - زاهدان - ایران -

snoferesti@ece.usb.ac.ir

^۲ گروه مهندسی کامپیوتر - دانشکده مهندسی برق و کامپیوتر - دانشگاه سیستان و بلوچستان - زاهدان - ایران -

mahshidmiri.ac@gmail.com

چکیده: تحلیل احساسات یکی از شاخه‌های مهم پردازش زبان طبیعی است که هدف آن طبقه‌بندی متون بر اساس احساس و نگرش نویسنده متن است. در زبان فارسی، متون نوشته شده در شبکه‌های اجتماعی غالباً کوتاه، بدون ساختار و مملو از عبارات محاوره‌ای و غیررسمی هستند که این ویژگی‌ها باعث می‌شود کارایی الگوریتم‌های تحلیل احساسات به طور چشمگیری کاهش یابد. هدف این مقاله ارائه روشی مبتنی بر یادگیری عمیق و واژه‌نامه حسی برای تحلیل احساسات متون فارسی نوشته شده در شبکه‌های اجتماعی است. به دلیل این که اغلب واژه‌نامه‌های حسی موجود در زبان فارسی از لحاظ اندازه کوچک و فاقد عبارات محاوره‌ای و غیررسمی هستند، ابتدا روشی برای گسترش واژه‌نامه‌های حسی موجود با افزودن عبارات محاوره‌ای پرکاربرد در رسانه‌های اجتماعی که به کمک ChatGPT تعیین قطبیت شده‌اند ارائه می‌شود. سپس از ترکیب واژه‌نامه حسی و شبکه عصبی پیچشی دو کاناله برای تعیین قطبیت متون استفاده می‌شود. نتایج ارزیابی‌های انجام گرفته نشان می‌دهد که با گسترش واژه‌نامه‌های حسی موجود با دو روش پیشنهادی، صحت الگوریتم تحلیل احساسات به ترتیب $1/74$ و $2/14$ درصد افزایش می‌یابد که نشان‌دهنده موفقیت ChatGPT در تعیین قطبیت عبارات محاوره‌ای فارسی است. همچنین، بکارگیری ویژگی‌های مستخرج از واژه‌نامه حسی در یک شبکه عصبی پیچشی دوکاناله منجر به افزایش دقت دو مدل پایه موردبررسی به میزان $1/6$ و $3/2$ درصد می‌شود.

واژه‌های کلیدی: تحلیل احساسات، تعیین قطبیت، یادگیری عمیق، شبکه عصبی پیچشی دوکاناله، واژه‌نامه حسی، عبارات محاوره‌ای

* نویسنده مسئول، snoferesti@ece.usb.ac.ir

Proposing an approach based on deep learning and sentiment lexicon for Persian sentiment analysis

Samira Noferesti ^{1*}, Associate Professor, Mahshid Miri ², Graduated Student

¹ Dept. of Information Technology Engineering, Faculty of Electrical and Computer Engineering, University of Sistan and Baluchestan, Zahedan, Iran, snoferesti@ece.usb.ac.ir

² Dept. of Computer Engineering, Faculty of Electrical and Computer Engineering, University of Sistan and Baluchestan, Zahedan, Iran, mahshidmiri.ac@gmail.com

Abstract: Sentiment analysis is one of the important branches of natural language processing, which aims to classify texts with respect to the feelings and attitudes of the author of the text. In Persian, most of the available sentiment lexicons are small in size and lack slang expressions and informal words. These features significantly reduce the performance of sentiment analysis algorithms. This paper aims to present a method based on deep learning and sentiment lexicons for sentiment analysis of Persian texts written on social networks. Since most existing sentiment lexicons in Persian language are small in size and lack slang and informal expressions, first, two methods based on ChatGPT are proposed to expand the existing Persian sentiment lexicons by adding slang expressions that are widely used in social media. Then, the combination of the sentiment lexicon and dual-channel convolutional neural network (DC-CNN) is used to determine the polarity of texts. Experimental results show that by expanding the existing sentiment lexicons with the two proposed methods, the accuracy of the sentiment analysis algorithm increases by 1.74 and 2.14 percent, respectively, which indicates the success of ChatGPT in polarity classification of Persian slang expressions. Also, employing the features extracted from the sentiment lexicon in a DC-CNN leads to an increase in the precision of the two base models by 1.6 and 3.2 percent.

Keywords: Sentiment analysis, Polarity classification, Deep learning, Dual-channel CNN, Sentiment lexicon

* Corresponding author

۱. مقدمه

تمایل افراد به نوشتن نظرات و دیدگاه‌های خود در بستر شبکه‌های اجتماعی، فرصت‌هایی را برای صاحبان کسب و کار، سیاستمداران و مدیران سازمان‌ها فراهم ساخته است که بازخوردهای ارزشمندی را از کاربران دریافت کنند که می‌تواند در شناسایی نقاط ضعف و قوت و بهبود تصمیم‌گیری‌های مدیریتی بسیار موثر باشد. به دلیل حجم عظیم نظرات موجود در وب و رشد روز افزون آنها، تحلیل دستی نظرات دشوار و گاه غیر ممکن است. تحلیل احساسات یکی از شاخه‌های مهم پردازش زبان طبیعی است که هدف آن طبقه‌بندی خودکار متون در کلاس‌های حسی از قبل تعیین شده مانند مثبت، منفی و خنثی می‌باشد. این شاخه علمی در سال‌های اخیر محبوبیت زیادی یافته و در دامنه‌های گوناگون از جمله سیاست، پزشکی و تجارت به کار گرفته شده است [۱]، [۲].

روش‌های موجود برای تحلیل احساسات را می‌توان به دو دسته کلی مبتنی بر واژه‌نامه و یادگیری ماشین تقسیم کرد [۳]. ایده اصلی روش‌های مبتنی بر واژه‌نامه، تبدیل متن به کیسه‌ای از کلمات، تعیین قطبیت هر کلمه بر اساس یک واژه‌نامه حسی و تعیین قطبیت کلی متن بر اساس قطبیت کلمات آن است. از این رو، کارایی این روش‌ها به شدت به میزان پوشش و کیفیت واژه‌نامه حسی مورد استفاده وابسته است. در مقابل، روش‌های یادگیری ماشین، با آموزش یک مدل بر روی پیکره‌ای از متون دارای برچسب قطبیت به تعیین قطبیت متون جدید می‌پردازند. کارایی روش‌های یادگیری ماشین تحت تاثیر کمیت و کیفیت پیکره آموزش و همچنین ویژگی‌های انتخاب شده برای آموزش مدل است. برخی از روش‌های یادگیری ماشین موجود از ویژگی‌های مبتنی بر واژه‌نامه حسی مانند حضور واژه‌های دارای قطبیت در متن، امتیاز قطبیت واژه‌ها، قطبیت کلی جمله بر اساس

واژه‌نامه حسی و تعداد صفات مثبت و منفی استفاده می‌کنند [۴] - [۶]. بنابراین کیفیت واژه‌نامه‌های حسی بر کارایی این دسته از روش‌های یادگیری ماشین نیز موثر است.

در سال‌های اخیر، مدل‌های یادگیری عمیق برای تحلیل احساسات مورد توجه محققان قرار گرفته‌اند که در مقایسه با روش‌های سنتی به دقت بالاتری در طبقه‌بندی احساسات دست یافته‌اند [۷]. علی‌رغم موفقیت‌های الگوریتم‌های یادگیری عمیق در طبقه‌بندی احساسات متون رسمی، به دلیل این که این الگوریتم‌ها اغلب برای جاسازی کلمات^۱ از پیکره‌های رسمی مانند ویکی‌پدیا استفاده می‌کنند، در تعیین قطبیت متون غیررسمی و محاوره‌ای از دقت کافی برخوردار نیستند. برای مقابله با چالش مذکور، در این مقاله از شبکه عصبی پیچشی دوکاناله (DC-^۲CNN) برای تحلیل احساسات متون فارسی نوشته شده در شبکه - های اجتماعی استفاده می‌شود. در این شبکه، ورودی یک کانال متن اولیه و ورودی کانال دیگر دنباله برچسب قطبیت عبارات متن است که از یک واژه‌نامه حسی استخراج می‌شود. بنابراین نیاز به یک واژه‌نامه حسی است که علاوه بر عبارات رسمی، عبارات محاوره‌ای رایج در زبان فارسی را نیز شامل شود.

در سال‌های اخیر، تلاش‌های متعددی جهت ساخت دستی و خودکار واژه‌نامه‌های حسی انجام شده است. برای نمونه، در زبان فارسی می‌توان به واژه‌نامه‌های حسی PerSent [۶]، حس‌نگار [۸] و LexiPers [۹] اشاره کرد. اغلب واژه‌نامه‌های موجود فاقد عبارات محاوره‌ای و غیررسمی هستند. این در حالی است که نظرات کاربران در شبکه‌های اجتماعی سرشار از عبارات غیررسمی است و عدم توجه به آنها باعث کاهش دقت الگوریتم‌های تحلیل احساسات می‌شود. بنابراین غنی‌سازی واژه‌نامه‌های موجود با عبارات محاوره‌ای و غیررسمی یکی از ضروریاتی است که باید به آن پرداخته شود. در این راستا، روشی

^۲ Dual-channel CNN

^۱ Word embedding

برای ساخت یک واژه‌نامه حسی از عبارات محاوره‌ای زبان فارسی با استفاده از ChatGPT³ پیشنهاد می‌شود. ChatGPT محصول جدید OpenAI⁴ و یک چت‌بات هوش مصنوعی مبتنی بر مدل‌های زبانی بزرگ⁴ است که در زمینه‌های متعدد از جمله پردازش زبان طبیعی مهارت دارد.

به صورت خلاصه، نوآوری‌های مقاله حاضر عبارتند از:

- ساخت خودکار یک واژه‌نامه حسی از عبارات محاوره‌ای و غیررسمی در زبان فارسی با کمک ChatGPT
- ارائه یک شبکه عصبی پیچشی دوکاناله که علاوه بر متن اولیه از دنباله برچسب‌های حسی عبارات متن نیز برای طبقه‌بندی احساسات بهره می‌برد.

ادامه مقاله به صورت زیر سازمان‌دهی شده است. در بخش ۲، تحقیقات پیشین در زمینه تحلیل احساسات معرفی می‌شود. بخش ۳ جزئیات روش پیشنهادی را شرح می‌دهد. در بخش ۴، به ارزیابی کارایی روش پیشنهادی و مقایسه آن با روش‌های موجود پرداخته می‌شود. در پایان، بخش ۵ نتیجه‌گیری می‌باشد.

۲. کارهای گذشته

در این بخش ابتدا به معرفی روش‌های موجود برای ساخت واژه‌نامه‌های حسی و سپس به مرور تحقیقات موجود در زمینه تحلیل احساسات با یادگیری عمیق به طور خاص برای زبان فارسی می‌پردازیم.

۱.۲. مرور روش‌های موجود برای ساخت واژه‌نامه‌های حسی

روش‌های موجود برای ساخت واژه‌نامه‌های حسی را می‌توان به چهار دسته کلی تقسیم کرد: روش‌های دستی، روش‌های مبتنی بر لغت‌نامه، روش‌های مبتنی بر پیکره و روش‌های ترکیبی.

تلاش‌های اولیه در جهت ساخت دستی واژه‌نامه‌های حسی بوده است. در این روش، یک گروه از افراد خبره به صورت دستی به واژه‌ها برچسب قطبیت تخصیص می‌دهند و در مواردی که اختلاف نظر وجود دارد، اغلب از رای اکثریت برای تعیین برچسب نهایی استفاده می‌شود. از جمله تحقیقاتی که به ساخت دستی واژه‌نامه‌های حسی در زبان فارسی پرداخته‌اند می‌توان به مراجع [۶، ۱۲-۱۰] اشاره کرد. ساخت دستی واژه‌نامه دشوار، هزینه‌بر و زمان‌بر است و به همین دلیل اغلب واژه‌نامه‌های حسی که به صورت دستی ساخته شده‌اند کوچک هستند و پوشش کمی دارند.

دسته دوم، روش‌های مبتنی بر لغت‌نامه هستند که اغلب با یک مجموعه کوچک از لغاتی که به صورت دستی تعیین قطبیت شده‌اند شروع کرده و با استفاده از شباهت معنایی کلمات و روابط نظیر مترادف و متضاد به تعیین قطبیت سایر لغات می‌پردازند. در زبان انگلیسی، معروف‌ترین و پرکاربردترین واژه‌نامه‌های این دسته SentiWordNet [۱۳] و SenticNet [۱۴] است.

در [۱۵] از نگاهت WordNet و FarsNet و الگوریتم قدم زدن تصادفی برای ساخت یک واژه‌نامه حسی فارسی استفاده شده است. ابتدا یک مجموعه اولیه از لغات انگلیسی با قطبیت مشخص تهیه شده و با استفاده از الگوریتم قدم زدن تصادفی قطبیت سایر لغات در شبکه واژه‌ها مشخص شده است. سپس با استفاده از پیوند بین مجموعه لغات هم‌معنی^۵ در نسخه اول FarsNet با شبکه واژه‌های انگلیسی و نیز مترجم گوگل برای گروه‌های هم‌معنی که پیوندی با FarsNet ندارند، یک واژه‌نامه حسی فارسی تهیه شده

⁵ Synset

³ <https://chat.openai.com/>

⁴ Large Language Model

است. در پایان نیز با کمک روابط مترادف و متضاد در شبکه واژه‌های فارسی، قطبیت کلمات اصلاح شده است.

در [۸] برای ساخت واژه‌نامه‌ای به نام حس‌نگار، ابتدا با نگاشت مجموعه کلمات هم‌معنی در WordNet به زبان فارسی، یک شبکه واژه به نام فردوس‌نت ایجاد شده است. سپس امتیاز حسی محاسبه شده برای هر مجموعه کلمات هم‌معنی در WordNet به متناظر آن در حس‌نگار نگاشت می‌شود.

روش‌های مبتنی بر پیکره، از روش‌های آماری و قواعد زبانی برای پیدا کردن قطبیت کلمات بر اساس هم‌وقوعی آنها با کلمات دارای قطبیت مشخص در یک پیکره استفاده می‌کنند [۱۶]. در [۱۷] ابتدا با کمک منابع انگلیسی و با استفاده از روش ترجمه، مجموعه‌ای از واژه‌های فارسی تهیه و به صورت دستی با مقادیر مثبت، منفی و خنثی برچسب‌گذاری شده است. سپس، طبقه‌بند یادگیری ماشین قطبیت واژه‌های فارسی را بر اساس امتیاز قطبیت معادل انگلیسی آنها در چهار منبع SentiWordNet، SenticNet، NRC و واژه‌نامه Liu تعیین می‌کند.

در [۹] یک روش ترکیبی پیشنهاد شده است که از هستان‌شناسی FarsNet برای ساخت یک واژه‌نامه حسی به نام LexiPers استفاده می‌کند. به این صورت که ابتدا یک لیست اولیه از لغات FarsNet به صورت دستی تعیین قطبیت شده است. سپس با کمک روابط بین کلمات و یک تابع مبتنی بر PMI^6 به گسترش واژه‌نامه حسی پرداخته شده است. در پایان، از واژه‌نامه حسی بدست آمده در مرحله قبل به عنوان مجموعه آموزش طبقه‌بندهای یادگیری ماشین استفاده شده است تا سایر مجموعه کلمات هم‌معنی در FarsNet تعیین قطبیت شوند.

همان‌طور که گفته شد، تحقیقات پیشین برای ساخت واژه‌نامه‌های حسی از فرهنگ‌لغت یا پیکره‌های متنی رسمی استفاده

می‌کنند. به همین دلیل، واژه‌نامه‌های موجود قادر به تعیین قطبیت بسیاری از نظرات ارزشمندی که توسط کاربران عمومی وب نوشته می‌شوند و به صورت محاوره‌ای هستند، نمی‌باشند. از این رو، گسترش واژه‌نامه‌های حسی موجود با اصطلاحات عامیانه و عبارات محاوره‌ای و غیررسمی امری ضروری است.

۲.۲. مرور روش‌های موجود در زمینه تحلیل احساسات

با یادگیری عمیق

در این بخش به طور خاص به بررسی تحلیل احساسات متون فارسی با یادگیری عمیق می‌پردازیم. در [۱۸]، دو مدل یادگیری عمیق مبتنی بر CNN و LSTM^۷ برای طبقه‌بندی نظرات کاربران فارسی زبان در دو دامنه فیلم و هتل پیشنهاد شده است. در این روش، بر روی متن ورودی پیش‌پردازش‌هایی مانند قطعه‌بندی، نرمال‌سازی و ریشه‌یابی انجام شده است. همچنین برای جاسازی کلمات از fastText [۱۹] استفاده شده است. نتایج ارزیابی‌های انجام گرفته نشان می‌دهد که bi-LSTM در دامنه فیلم و CNN در دامنه هتل عملکرد بهتری داشته‌اند.

در [۲۰]، یک روش مبتنی بر یادگیری تطبیقی برای تحلیل احساسات پیشنهاد شده است. در مرحله اول این روش، تطابق بین ویژگی‌های دامنه‌های مختلف با ویژگی‌های محوری تعیین می‌شود. ویژگی‌های محوری، ویژگی‌هایی هستند که در یادگیری در دو دامنه مختلف رفتار یکسانی دارند و می‌توانند برای هر دو دامنه استفاده شوند. سپس از این ویژگی‌های محوری برای استخراج ویژگی‌های غیرمحوری و مستقل استفاده می‌شود. در مرحله دوم با آموزش یک مدل شبکه عصبی پیچشی برای روی فضای ویژگی استخراج شده در گام قبل قطبیت متون تعیین می‌شود.

^۷ Long Short-Term Memory

^۶ Pointwise Mutual Information

در [۲۱]، معماری جدیدی متشکل از ParsBERT و bi-LSTM برای طبقه‌بندی نظرات کاربران دیجی‌کالا ارائه شده است. در ابتدا با استفاده از BERT، متن ورودی به سه بردار جاسازی توکن‌ها، قطعات و مکان رمزگذاری می‌شود و این بردارها به عنوان ورودی به ParsBERT داده می‌شوند. خروجی ParsBERT یک تانسور است که پس از اعمال حذف تصادفی^۸ بر روی آن به عنوان ورودی به bi-LSTM داده می‌شود که وظیفه طبقه‌بندی نظرات را بر عهده دارد.

در [۲۲]، برای تحلیل احساسات توئیتهای سیاسی فارسی زبانان، الگوریتم‌های یادگیری ماشین متعددی شامل جنگل تصادفی، ماشین بردار پشتیبان، شبکه‌های عصبی و همچنین مدل‌های یادگیری عمیق با روش‌های مختلف جاسازی کلمات با Word2Vec، fastText و ParsBERT بکار گرفته شده است که بهترین نتیجه توسط مدل CNN+biLSTM با جاسازی ParsBERT حاصل شده است.

در روش پیشنهادی [۲۳]، ابتدا بر روی مجموعه آموزش داده‌افزایی انجام شده است. سپس پیش‌پردازش‌های حذف ایست-واژه‌ها^۹، نرمال‌سازی و تقطیع واژه‌ها بر روی متن ورودی انجام گرفته است و در پایان نظرات کاربران با کمک طبقه‌بند CNN تعیین قطبیت شده‌اند. طبقه‌بند CNN پیشنهادی دارای لایه‌های اصلی جاسازی کلمات، پیچشی (با ۱۲۸ فیلتر با اندازه هسته ۳)، ادغام حداکثری^{۱۰} و کاملاً متصل^{۱۱} با تابع بیشینه نرم است. همچنین، ابعاد بردار کلمات ورودی برابر ۱۰۰، نرخ حذف تصادفی در لایه ماقبل آخر برابر ۰/۵ و الگوریتم بهینه‌سازی آدام در نظر گرفته شده است.

در [۲۴]، از چهار مدل یادگیری عمیق به نام‌های CNN، GRU، LSTM و bi-LSTM برای تحلیل احساسات نظرات کاربران

فارسی زبان در دامنه فیلم استفاده شده است که بر اساس ارزیابی‌های انجام گرفته، CNN در مقایسه با سایر مدل‌ها به نتایج بهتری دست یافته است. معماری CNN پیشنهادی به ترتیب شامل لایه جاسازی کلمات، حذف تصادفی با نرخ ۰/۵، لایه پیچشی با ۲۵۶ فیلتر با اندازه ۳، لایه ادغام حداکثری سراسری و دو لایه کاملاً متصل با توابع فعالساز ReLU و سیگموئید است.

در [۲۵]، دو معماری یادگیری عمیق برای طبقه‌بندی جملات به دو دسته مثبت و منفی پیشنهاد شده است که بهترین نتایج با مدل مبتنی bi-LSTM حاصل شده است. معماری این مدل متشکل از لایه‌های جاسازی کلمات، bi-LSTM، ادغام حداکثری سراسری، حذف با نرخ ۰/۲، کاملاً متصل با تابع فعالساز ReLU، حذف با نرخ ۰/۱ و کاملاً متصل با تابع فعالساز سیگموئید است.

در [۲۶]، از یادگیری جمعی برای تحلیل احساسات نظرات کاربران اینستاگرام استفاده شده است. روش پیشنهادی این پژوهش سه مرحله دارد: در مرحله اول، پیش‌پردازش و جاسازی کلمات انجام می‌شود. در مرحله دوم، با استفاده از چهار مدل یادگیری عمیق شامل CNN، LSTM، CNN-LSTM و LSTM-CNN با تعداد دوره‌های اجرای کم به طبقه‌بندی احساسات پرداخته می‌شود. در مرحله سوم، با دو روش پرسپترون چندلایه و رای‌گیری مبتنی بر وزن (وزن یک مدل، دقت طبقه‌بندی آن مدل است)، نتایج چهار مدل تجمیع می‌گردد. نتایج ارزیابی‌های انجام گرفته در این پژوهش نشان می‌دهد که تجمیع نتایج چهار مدل با روش رای‌گیری به بیشترین دقت در طبقه‌بندی احساسات مجموعه تست رسیده است.

در سال‌های اخیر، مدل‌های مبتنی بر ترنسفورمر برای تحلیل احساسات متون فارسی مورد توجه قرار گرفته‌اند. در [۲۷]، ParsBERT برای وظیفه تحلیل احساسات تنظیم شده است. ParsBERT یک مدل زبانی پیشرفته است که بر اساس معماری

¹⁰ Max-Pooling

¹¹ Fully connected

⁸ Dropout

⁹ Stop words

BERT طراحی شده و بر روی مجموعه داده های فارسی متعددی از جمله ویکی پدیا، دیجی کالا، الی گشت، چطور و غیره آموزش داده شده است [۲۸]. این مدل از تکنیک های یادگیری عمیق و معماری ترنسفورمر بهره می برد و می تواند به طور مؤثری برای وظایف مختلف پردازش زبان طبیعی از جمله تحلیل احساسات تنظیم شود. در معماری پیشنهادی [۲۷]، ابتدا متن ورودی توسط رمزگذار BERT به سه نوع بردار جاسازی (توکن، بخش و موقعیت) تبدیل می شود. سپس این بردارها به ParsBERT داده می شوند و خروجی ParsBERT به ترتیب از لایه های حذف تصادفی، bi-LSTM و مسطح سازی و طبقه بندی با تابع فعال سازی سیگموئید عبور می کند.

در [۲۹] نیز روشی مبتنی بر ParsBERT برای تحلیل احساسات مبتنی بر جنبه پیشنهاد شده است. در این روش، مدل پیش آموخته ParsBERT بر روی یک مجموعه داده حاوی برچسب قطبیت از نظرات کاربران دیجی کالا تنظیم شده است. همچنین برای تشخیص دقیق تر جنبه ها، قبل از تزریق ورودی به مدل، غنی سازی معنایی جنبه ها با کمک استخراج کلمات مترادف هر جنبه از فرهنگ لغات معین و دهخدا انجام شده است.

در [۳۰]، با پیش آموزش یک مدل مبتنی بر BERT بر روی پیکره وبلاگ های فارسی زبان مدل FaBERT ایجاد شده است که تنظیم آن برای وظیفه تحلیل احساسات در مقایسه با ParsBERT بر روی مجموعه داده های نظرات کاربران دیجی کالا و توییتهای فارسی به دقت بالاتری دست یافته است.

بر اساس تحقیقات نویسندگان، تابحال برای تحلیل احساسات متون فارسی از شبکه های عصبی چندکاناله استفاده نشده است. واژه نامه های حسی اطلاعات ارزشمندی را برای تحلیل احساسات فراهم می کنند که بکارگیری آنها در مدل های یادگیری عمیق می تواند کمک بسیاری به درک قطبیت درست متن نماید. در این راستا، در مقاله حاضر، ویژگی های مستخرج از واژه نامه حسی به

همراه متن ورودی به یک شبکه عصبی پیچشی دو کاناله برای تحلیل احساسات داده می شود که جزئیات آن در بخش بعد شرح داده می شود.

۳. روش پیشنهادی

در این بخش روش پیشنهادی برای تحلیل احساسات متون غیررسمی در زبان فارسی تشریح می گردد. روشی پیشنهادی ترکیبی از روش های مبتنی بر واژه نامه و یادگیری عمیق است. یکی از رایج ترین مدل های یادگیری عمیق در تحلیل احساسات متون فارسی، شبکه های عصبی پیچشی هستند که برای استخراج ویژگی های زمینه از جاسازی کلمات مستخرج از پیکره های رسمی مانند ویکی پدیا استفاده می کنند و به همین دلیل در طبقه بندی متون عامیانه کارایی لازم را ندارند. در این مقاله، برای افزایش کارایی شبکه عصبی پیچشی در تحلیل احساسات متون غیررسمی، یک شبکه دو کاناله در نظر گرفته شده است که ورودی یک کانال متن نظر کاربر و ورودی کانال دیگر دنباله ای از برچسب قطبیت کلمات متن است که با کمک یک واژه نامه حسی شامل کلمات رسمی و عبارات محاوره ای رایج در زبان فارسی تعیین قطبیت شده اند.

در ادامه این بخش، ابتدا روش پیشنهادی برای ساخت یک واژه نامه حسی از عبارات محاوره ای فارسی معرفی می گردد و سپس جزئیات روش پیشنهادی برای طبقه بندی احساسات شرح داده می شود.

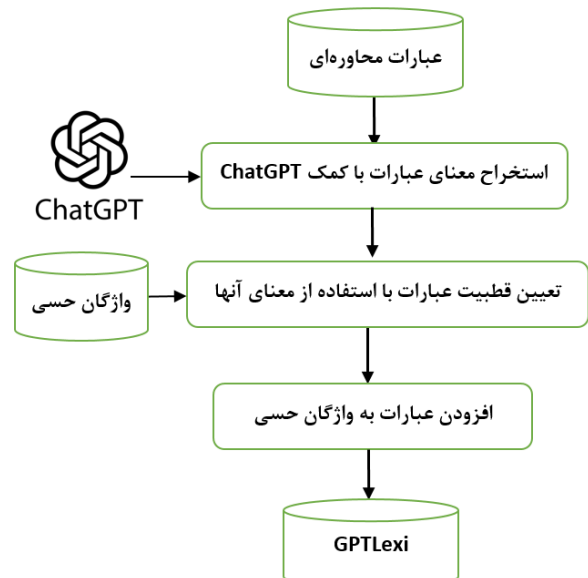
۳.۱. روش پیشنهادی برای ساخت واژه نامه حسی

هدف روش پیشنهادی، گسترش واژه نامه های حسی زبان فارسی با عبارات محاوره ای و غیررسمی است. برای این منظور، ابتدا یک مجموعه از عبارات محاوره ای و اصطلاحات غیررسمی رایج در نوشتار شبکه های اجتماعی جمع آوری شده است. برای جمع آوری

عبارات، از مجموعه داده عبارات محاوره‌ای زبان فارسی [۳۱] و نیز جستجو در google کمک گرفته شده است. بدین ترتیب لیستی شامل ۳۵۰ عبارت محاوره‌ای فارسی تهیه شده است. سپس به تعیین قطبیت این عبارات با ChatGPT پرداخته شده است. برای این منظور، دو روش پیشنهاد شده است که در ادامه جزئیات آنها شرح داده می‌شود.

۱.۱.۳. روش پیشنهادی برای ساخت GPTLexi

شکل (۱) مراحل روش پیشنهادی اول برای ساخت یک واژه‌نامه حسی از عبارات محاوره‌ای فارسی به نام GPTLexi را نشان می‌دهد. مطابق شکل (۱)، برای هر عبارت محاوره‌ای، ابتدا معنای آن با کمک ChatGPT بدست می‌آید و سپس پاسخ ChatGPT (مفهوم عبارت محاوره‌ای) با کمک یک واژه‌نامه حسی تعیین قطبیت می‌شود. دلیل این امر این است که واژه‌نامه‌های حسی موجود فاقد عبارات محاوره‌ای رایج در شبکه‌های اجتماعی هستند ولی پاسخی که ChatGPT می‌دهد یک متن ادبی است که انتظار می‌رود بتوان لغات آن را با واژه‌نامه حسی تعیین قطبیت کرد.



شکل (۱): مراحل ساخت واژه‌نامه حسی GPTLexi

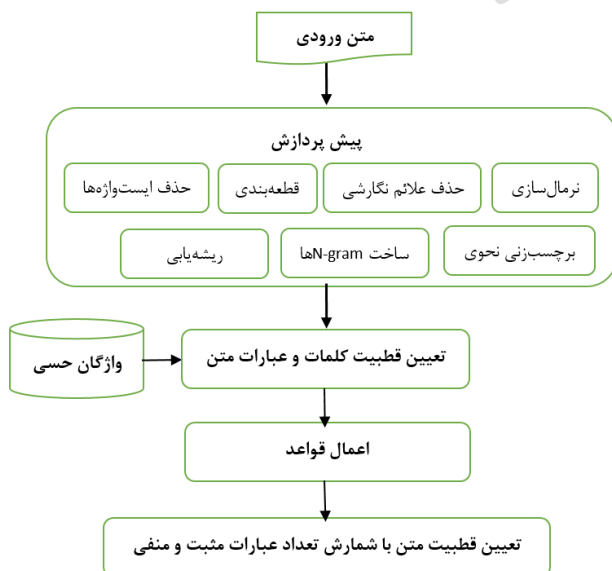
برای مثال، برای عبارت محاوره‌ای "آش‌ولاش" سوالی که از ChatGPT پرسیده شده و پاسخ آن به صورت زیر است:

- سوال: عبارت عامیانه "آش‌ولاش" در فارسی یعنی چه؟

- پاسخ: عبارت "آش‌ولاش" در زبان فارسی به معنای غیرترتیبی، به هم ریخته، یا در هر جور شکلی به نظر کثیف یا غیرسازمان‌یافته آمده باشد، اشاره دارد. این اصطلاح به طور عام برای توصیف چیزی که به نظر نظم و ترتیب ندارد، به کار می‌رود.

در مثال فوق، لغاتی مانند "کثیف" که در توصیف کلمه "آش‌ولاش" آمده است را می‌توان با واژه‌نامه‌های موجود تعیین قطبیت کرد، هرچند خود کلمه "آش‌ولاش" در واژه‌نامه حسی نباشد.

برای تعیین قطبیت پاسخ داده شده توسط ChatGPT از روش شمارش کلمات مثبت و منفی به همراه مجموعه‌ای از قواعد برای مدیریت معکوس‌کننده‌ها و اصلاح قطبیت استفاده می‌شود. مراحل روش پیشنهادی برای تعیین قطبیت مبتنی بر واژه‌نامه در شکل (۲) نشان داده شده است.



شکل (۲): روش پیشنهادی برای تعیین قطبیت مبتنی بر واژه‌نامه

مطابق شکل (۲)، اولین مرحله در تعیین قطبیت متن، پیش‌پردازش است. پیش‌پردازش متن شامل ۷ مرحله نرمال‌سازی،

حذف علائم نگارشی، قطعه‌بندی، حذف ایست‌واژه‌ها، برچسب زنی نحوی، ساخت عبارات و ریشه‌یابی است. نرمال‌سازی با هدف یکسان‌سازی شکل‌های مختلف کلمات با جایگزین کردن حروف استاندارد در متن ورودی انجام می‌شود. در مرحله بعد علائم نگارشی مانند !، ؟ و . حذف می‌گردند. قطعه‌بندی، کلمات متن را مشخص می‌کند. سپس ایست‌واژه‌های متن با کمک لیستی شامل ۶۰۵ ایست‌واژه در زبان فارسی، شناسایی و حذف می‌شوند. برچسب‌زن نحوی، نقش نحوی کلمات را مشخص می‌سازد که در قواعد مدیریت معکوس‌کننده‌ها در شناسایی افعال منفی استفاده می‌شود. در پایان نیز ترکیب دوتایی^{۱۲} و سه‌تایی^{۱۳} کلمات برای تعیین عبارات متن تشکیل می‌شوند و کلیه کلمات و عبارات متن ریشه‌یابی می‌شوند. دلیل ریشه‌یابی این است که بسیاری از کلمات و عبارات موجود در واژه‌نامه‌های حسی به شکل ریشه خود هستند.

پس از پیش‌پردازش، ابتدا به جستجوی عبارات سه کلمه‌ای متن و ریشه آنها در واژه‌نامه حسی پرداخته می‌شود. سپس عبارات دو کلمه‌ای و ریشه آنها و در پایان نیز، کلمات تکی و ریشه آنها جستجو می‌شوند. پس از این که به کمک واژه‌نامه حسی، عبارات مثبت و منفی متن تعیین شدند، قواعد کمکی زیر برای تعیین قطبیت و مدیریت معکوس‌کننده‌ها استفاده می‌شوند:

- قاعده ۱: اگر در متن "و" یا "یا" مشاهده شد و کلمه قبل و بعد آن برچسب نحوی یکسان داشتند و یکی از این واژه‌ها در واژه‌نامه حسی وجود داشت و دیگری نه، قطبیت هر دو واژه یکسان در نظر گرفته می‌شود. اگر کلمه بعد از "و" یا "یا" برچسب قید داشت ولی دو کلمه بعد از آن برچسب نحوی یکسانی با کلمه قبلی داشت نیز این قاعده اعمال

می‌شود. در جدول (۱)، ردیف ۱ دو مثال از این قاعده را نشان می‌دهند.

- قاعده ۲: اگر کلمه‌ای با پیشوند "بی" در واژه‌نامه قطبیت منفی داشته باشد، این کلمه با پیشوند "با" قطبیت مثبت خواهد داشت و برعکس. در جدول (۱)، ردیف ۲ مثال‌هایی از این قاعده را نشان می‌دهد.

- قاعده ۳: اگر قبل از یک عبارت مثبت یک کلمه معکوس‌کننده مانند "کمبود"، "فقدان"، "نبود"، "غیر"، "عدم" و "نه" باشد، حس عبارت معکوس می‌شود. در جدول (۱)، ردیف ۳ دو مثال برای این قاعده را نشان می‌دهد. در واژه‌نامه حسی، کلمه "محبت" قطبیت مثبت دارد ولی به دلیل کلمه "کمبود" که قبل از آن قرار گرفته است، قطبیت کلمه "محبت" معکوس می‌شود. به طور مشابه، در مثال دوم، حضور کلمه "نه" از کلمه منفی "چاق" باعث معکوس شدن قطبیت آن می‌شود.

- قاعده ۴: اگر پس از یک عبارت اسمی مثبت، لغت "کم" یا "کمی" بیاید، قطبیت عبارت معکوس می‌شود. در ردیف ۴ جدول (۱)، مثالی برای این قاعده ذکر شده است. در این مثال، با وجود این که قطبیت کلمه "درخشش" مثبت است، حضور کلمه "کمی" بعد از آن باعث منفی شدن قطبیت جمله می‌شود.

- قاعده ۵: اگر تا سه کلمه بعد از عبارت حسی یک فعل منفی وجود داشته باشد، قطبیت عبارت معکوس می‌شود. برای شناسایی افعال منفی به این صورت عمل شده است که اگر ابتدای فعل حرف "ن" باشد و پس از حذف آن همچنان یک کلمه معتبر با برچسب نحوی فعل حاصل شود این کلمه به عنوان یک فعل منفی در نظر گرفته می‌شود. همچنین

در روش دوم، برای هر عبارت محاوره‌ای، قطبیت آن از ChatGPT پرسیده می‌شود. برای مثال، برای عبارت "آش‌ولاش" سوال زیر از ChatGPT پرسیده شده است:

- سوال: حس عبارت عامیانه "آش‌ولاش" در فارسی چیست؟
مثبت، منفی یا خنثی؟

- پاسخ: عبارت "آش‌ولاش" در فارسی یک عبارت عامیانه است که معنای منفی دارد. این عبارت معمولاً برای توصیف وضعیت یا حالت چیزی یا کسی که به شدت آسیب دیده، خراب یا به هم ریخته است، به کار می‌رود.

بر اساس پاسخ ChatGPT، قطبیت عبارت محاوره‌ای (برای مثال قطبیت منفی برای عبارت فوق) تعیین می‌شود. با این وجود، برای برخی عبارات پاسخ قطعی وجود ندارد. برای نمونه، در مثال زیر نمی‌توان بر اساس پاسخ ChatGPT به قطبیت واژه "سوگلی" پی برد.

- سوال: حس عبارت عامیانه "سوگلی" در فارسی چیست؟
مثبت، منفی یا خنثی؟

- پاسخ: عبارت "سوگلی" در فارسی ممکن است به دلایل مختلف در جوامع مختلف و در شرایط مختلف، حس متفاوتی داشته باشد. اما به طور کلی، استفاده از این اصطلاح می‌تواند منفی، مثبت یا حتی خنثی باشد، بسته به سیاق و موقعیتی که در آن استفاده می‌شود. بنابراین، تفسیر حس "سوگلی" در هر موقعیت به شرایط و سیاق خاصی که این اصطلاح در آن به کار رفته است، بستگی دارد.

در مواردی که ChatGPT برای یک عبارت چندین قطبیت مشخص کرده است، قطبیت خنثی در نظر گرفته شده است یا به بیانی دیگر از آن عبارت صرف‌نظر شده است. سایر عبارات با روش

لیستی از افعال منفی که از قاعده فوق پیروی نمی‌کنند تهیه شده است مانند "نیست" و "از دست داد" که در شناسایی افعال منفی مورد استفاده قرار می‌گیرند. ردیف ۵ در جدول (۱)، دو کاربرد این قاعده را نشان می‌دهد.

جدول (۱): مثال‌هایی از قواعد

شماره قاعده	مثال
۱	سوال: <u>خلایقانه</u> و <u>باحال</u> بود خلاصه بگم <u>بچه ننه</u> و <u>خیلی لوس</u>
۲	بی‌کلاس و با کلاس، بی‌ادب و با ادب او دچار <u>کمبود محبت</u> است.
۳	من نرمالم و نه چاق. <u>درخشش کمی</u> دارد.
۴	اونقدر <u>باحال</u> نبود. <u>بی‌وفا نیستی</u> .
۵	

پس از اعمال قواعد، تعداد عبارات مثبت و منفی شمرده می‌شود. اگر تعداد عبارات مثبت بیشتر باشد قطبیت کلی متن مثبت و اگر تعداد عبارات منفی بیشتر باشد قطبیت کلی متن منفی در نظر گرفته می‌شود. اگر تعداد عبارات مثبت و منفی متن برابر باشند، قطبیت خنثی به متن تخصیص داده می‌شود.

با توصیف فوق، مجموعه‌ای از عبارات محاوره‌ای دارای برچسب قطبیت ساخته می‌شود که آن‌ها را به واژه‌نامه حسی اضافه می‌کنیم. از آنجا که واژه‌نامه حسی اولیه تنها حاوی واژه‌های مثبت و منفی بود، از عبارات محاوره‌ای خنثی صرف‌نظر شده و تنها عبارات حاوی قطبیت به واژه‌نامه حسی اضافه شده‌اند. واژه‌نامه حسی حاصل از این روش را GPTLexi می‌نامیم.

۲.۱.۳. روش پیشنهادی برای ساخت SentiGPT

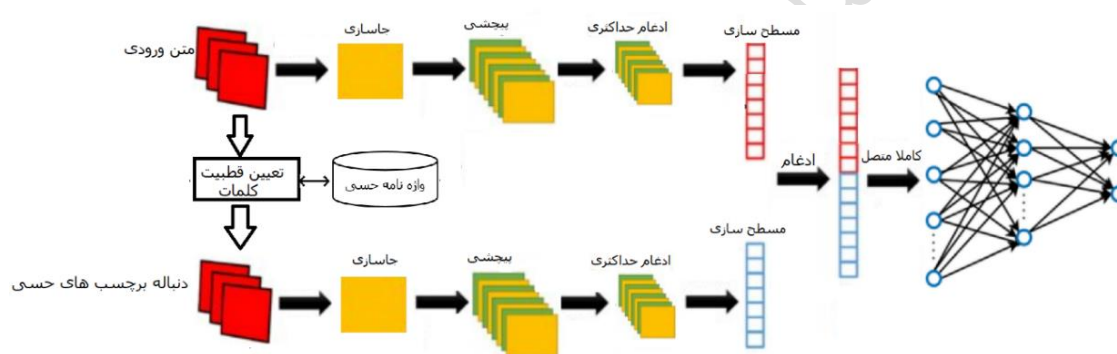
قطعه‌بندی، حذف علائم نگارشی، نرمال‌سازی و حذف ایست‌واژه-ها می‌شود.

مذکور تعیین قطبیت شده و با افزودن آن به واژه‌نامه‌های حسی موجود، واژه‌نامه‌ای ایجاد شده است که آن را SentiGPT می‌نامیم.

۲.۳. روش پیشنهادی برای تحلیل احساسات

هر یک از کانال‌ها در معماری پیشنهادی متشکل از لایه‌های جاسازی، پیچشی، ادغام حداکثری و مسطح‌سازی^{۱۴} است که می‌توان حذف و نرمال‌سازی دسته‌ای را نیز به آن افزود. بردارهای خروجی دو کانال ادغام می‌شوند و به عنوان ورودی به یک لایه کاملاً متصل داده می‌شود. آخرین لایه نیز یک لایه کاملاً متصل با تابع فعال‌ساز سیگموئید است که وظیفه طبقه‌بندی نظرات در دو دسته مثبت یا منفی را دارد.

روش پیشنهادی برای تحلیل احساسات یک شبکه عصبی پیچشی دوکاناله است که معماری آن در شکل (۳) نشان داده شده است. ورودی کانال اول، متن پیش‌پردازش شده و ورودی کانال دوم، دنباله‌ای از برچسب قطبیت هر یک از کلمات متن اولیه است که به کمک واژه‌نامه حسی بدست آمده است. پیش‌پردازش‌ها شامل



شکل (۳): معماری پیشنهادی برای تحلیل احساسات متون فارسی

۴. نتایج

برای پیاده‌سازی روش پیشنهادی از زبان برنامه‌نویسی پایتون نسخه 3.7.9 و ChatGPT 3.5 استفاده شده است. در این بخش، ابتدا به معرفی مجموعه داده‌ها و واژه‌نامه‌های حسی مورد استفاده و سپس به ارزیابی کارایی روش‌های پیشنهادی پرداخته می‌شود.

۱.۴. معرفی مجموعه داده‌ها و واژه‌نامه‌های حسی مورد

استفاده

مجموعه داده مورد استفاده در این مقاله، مجموعه داده نظرکاوی و تحلیل احساسات جمع‌آوری شده از توییترهای فارسی [۳۲] است که مشخصات آن در جدول (۲) گزارش شده است.

جدول (۲): مشخصات مجموعه داده

۱۲۶۱	تعداد نظرات
۴۸۹	تعداد نظرات مثبت
۷۷۲	تعداد نظرات منفی
۱۴/۹۵	متوسط طول نظرات
۵۵۷۱	تعداد کلمات متمایز
۱۸۸۸۵	تعداد کل کلمات

در روش پیشنهادی از ادغام سه واژه‌نامه حسی شناخته شده برای زبان فارسی به نام‌های LexiPers [۹]، PerSent [۶] و Amiri

¹⁴ Flatten

۱۳/۵ درصد از نظرات، حداقل یک عبارت محاوره‌ای که در واژه‌نامه‌های قبلی نبوده‌اند را در برداشته‌اند.

۲.۴. ارزیابی کارایی روش‌های پیشنهادی برای گسترش

واژه‌نامه‌های حسی

در جدول (۴)، کارایی تعیین قطبیت نظرات مجموعه داده توصیف شده در جدول (۲) توسط واژه‌نامه‌های LexiPers، PerSent و Amiri بر حسب معیار صحت^{۱۵} ارائه شده و با صحت حاصل از تعیین قطبیت توسط واژه‌نامه ترکیبی مقایسه شده است. منظور از صحت، درصد نظراتی است که توسط الگوریتم تعیین قطبیت به درستی برچسب زده شده‌اند. در همه موارد تعیین قطبیت متن با روش پیشنهادی که مراحل آن در شکل (۲) نشان داده شده است انجام شده است.

جدول (۴): صحت واژه‌نامه‌های حسی مختلف در تعیین قطبیت

نظرات

نام واژه‌نامه	صحت
LexiPers	۴۵/۲
PerSent	۲۳/۹۵
Amiri	۵۳/۱۳
واژه‌نامه ترکیبی	۵۷/۱۸

همان‌طور که در جدول (۴) مشاهده می‌شود، واژه‌نامه ترکیبی صحت بالاتری در تعیین قطبیت نظرات دارد. در مقایسه با واژه‌نامه Amiri که عملکرد بهتری نسبت به دو واژه‌نامه دیگر دارد، صحت تعیین قطبیت با واژه‌نامه ترکیبی حدود ۴ درصد بهبود داشته است.

با توجه به نتایج حاصل از جدول (۴)، در ادامه از واژه‌نامه ترکیبی استفاده شده که با دو روش پیشنهادی و به کمک ChatGPT گسترش داده شده است. همان‌طور که در بخش ۳

[۱۰] استفاده شده است که در بخش ۲ نحوه ساخت آنها توصیف شد. دلیل ادغام این سه واژه‌نامه، پوشش کم هر یک از واژه‌نامه‌ها به تنهایی و به تبع آن دقت کم آنها در تعیین قطبیت متون است. در هنگام ادغام سه واژه‌نامه، از عباراتی مانند "سبک"، "بلند" و "سرد" که در واژه‌نامه‌های مختلف قطبیت متفاوت دارند صرف‌نظر شده است. در واقع، قطبیت این لغات وابسته به زمینه است که در کار آتی به آن خواهیم پرداخت.

با ادغام سه واژه‌نامه حسی مذکور و با صرف‌نظر کردن از لغات خنثی، واژه‌نامه‌ای شامل ۴۹۱۸ عبارت دارای قطبیت حاصل شده است که در ادامه آن را "واژه‌نامه ترکیبی" می‌نامیم. مشخصات سه واژه‌نامه مذکور و واژه‌نامه ترکیبی در جدول (۳) آورده شده است. از آنجا که در واژه‌نامه PerSent، امتیاز قطبیت هر کلمه با عددی در بازه [-1,1] مشخص شده است از آستانه ۰/۲ برای تعیین واژه‌های مثبت و منفی استفاده شده است. به این صورت که اگر امتیاز کلمه‌ای بیشتر از ۰/۲ باشد مثبت و اگر امتیاز آن کمتر از ۰/۲- باشد منفی و در غیر این صورت خنثی در نظر گرفته شده است.

جدول (۳): مشخصات واژه‌نامه‌های حسی مورد استفاده

نام واژه‌نامه	تعداد عبارات مثبت	تعداد عبارات منفی
LexiPers	۹۹۵	۱۴۶۶
PerSent	۹۸	۱۰۹
Amiri	۱۲۹۳	۱۴۵۹
واژه‌نامه ترکیبی	۲۱۴۷	۲۷۷۱

بر اساس بررسی‌های انجام‌شده، ۸۰/۲۵ درصد از نظرات موجود در مجموعه داده، حداقل شامل یک واژه بوده‌اند که قطبیت آن توسط واژه‌نامه ترکیبی تعیین شده است. علاوه بر این، تقریباً

جدول (۶): کارایی قواعد پیشنهادی در بهبود صحت الگوریتم تعیین

قطبیت

نام واژه‌نامه	بدون قواعد	با قواعد
واژه‌نامه ترکیبی	۵۵/۹۹	۵۷/۱۸
GPTLexi	۵۷/۸۹	۵۸/۹۲
SentiGPT	۵۸/۴۵	۵۹/۳۲

با توجه به این که واژه‌نامه SentiGPT صحت بالاتری در تعیین قطبیت نظرات دارد، در ارزیابی‌های بعدی تنها از این واژه‌نامه استفاده شده است.

۳.۴. ارزیابی کارایی روش پیشنهادی برای تحلیل

احساسات

برای ارزیابی روش پیشنهادی نیاز به یک مجموعه تست است. در این راستا، به صورت تصادفی ۲۰ درصد از رکوردهای مجموعه داده توصیف شده در جدول (۲) به عنوان مجموعه تست و سایر رکوردها به عنوان مجموعه آموزش انتخاب شده‌اند.

روش پیشنهادی با در نظر گرفتن دو معماری مختلف که در مراجع [۲۳] و [۲۴] ارائه شده‌اند، مورد بررسی قرار گرفته است. به صورت دقیق‌تر، معماری هر یک از کانال‌ها بر اساس معماری پیشنهادی مراجع [۲۳] و [۲۴] تعریف شده است و بعد از آن، مطابق شکل (۳)، خروجی لایه مسطح‌سازی کانال‌ها ادغام شده و به لایه‌های کاملاً متصل داده می‌شود. خلاصه معماری‌های مورد استفاده و مقادیر مورد استفاده برای هایپرپارامترها در جداول (۷) و (۸) ارائه شده است.

توصیف شد، در روش پیشنهادی اول، یک واژه‌نامه حسی به نام GPTLexi و در روش پیشنهادی دوم یک واژه‌نامه حسی به نام SentiGPT ساخته شده است. در جدول (۵)، صحت این واژه‌نامه‌ها در تعیین قطبیت نظرات مجموعه داده گزارش شده است. در مقایسه با واژه‌نامه ترکیبی، واژه‌نامه‌های GPTLexi و SentiGPT به ترتیب صحت الگوریتم تعیین قطبیت را ۱/۷۴ و ۲/۱۴ درصد بهبود داده‌اند. دلیل اصلی این بهبود، وجود عبارات محاوره‌ای و غیررسمی در واژه‌نامه‌های ساخته شده توسط روش‌های پیشنهادی است که به کرات در نظرات کاربران در شبکه‌های اجتماعی مشاهده می‌شوند. همچنین انتظار می‌رود با بررسی تعداد عبارات محاوره‌ای بیشتر بهبود بیشتری در تعیین قطبیت نظرات حاصل شود.

جدول (۵): صحت واژه‌نامه‌های ساخته شده با روش‌های پیشنهادی

در تعیین قطبیت نظرات

نام واژه‌نامه	صحت
واژه‌نامه ترکیبی	۵۷/۱۸
GPTLexi	۵۸/۹۲
SentiGPT	۵۹/۳۲

همان‌طور که در بخش ۳ توصیف شد، برای تعیین قطبیت یک متن با کمک واژه‌نامه حسی، از یک مجموعه از قواعد از پیش تعریف شده استفاده می‌شود. جدول (۶) صحت الگوریتم تعیین قطبیت را در حالتی که در کنار واژه‌نامه حسی از این مجموعه قواعد نیز استفاده می‌شود با حالتی که قطبیت متن فقط با کمک واژه‌نامه حسی بدست می‌آید، مقایسه می‌کند. مطابق جدول (۶)، استفاده از قواعد معرفی شده در این مقاله، صحت واژه‌نامه‌های ترکیبی، GPTLexi و SentiGPT را به ترتیب ۱/۱۹، ۱/۰۳ و ۰/۸۷ درصد بهبود داده است.

جدول (۷): خلاصه معماری پیشنهادی اول

مشخصات هر یک از کانالها	
Ebmedding	کانال اول: Persian-Wikipedia-GloVe با ابعاد ۵۰
Convolutional	برای بردار کلمات
Dropout	شامل ۱۲۸ فیلتر با اندازه ۳ و تابع فعالساز ReLU
Max pooling	با نرخ ۰/۵
Flatten	با اندازه ۲
بعد از ادغام خروجی دو کانال	
Fully connected	فعالساز ReLU
Fully connected	فعالساز Sigmoid
مقادیر هایپر پارامترها	
بهینه ساز	Adam
تابع هزینه	Cross-Entropy Loss
نرخ یادگیری	۰/۰۰۰۵
اندازه دسته	۶۴
تعداد دورها	۲۰

جدول (۸): خلاصه معماری پیشنهادی دوم

مشخصات هر یک از کانالها	
Ebmedding	کانال اول: با ابعاد ۱۰۰ برای بردار کلمات
Dropout	با نرخ ۰/۵
Convolutional	شامل ۲۵۶ فیلتر با اندازه ۳ و تابع فعالساز ReLU
Batch Normalization	-
Global Max pooling	با اندازه ۲
Flatten	-
بعد از ادغام خروجی دو کانال	
Fully connected	فعالساز ReLU
Fully connected	فعالساز Sigmoid
مقادیر هایپر پارامترها	
بهینه ساز	Adam
تابع هزینه	Cross-Entropy Loss
نرخ یادگیری	۰/۰۰۰۵
اندازه دسته	۵۰
تعداد دورها	۲۰

در جدول (۹)، کارایی روش های پیشنهادی مبتنی بر شبکه عصبی پیچشی دوکاناله (DC-CNN) با روش های پایه معرفی شده در مراجع [۲۳] و [۲۴] که مبتنی بر شبکه عصبی پیچشی تک کاناله (1C-CNN) با معماری مشابه هستند و همچنین روش های مبتنی بر bi-LSTM [۲۵] و Voting-Ensemble [۲۶] مقایسه شده است. جزئیات روش های مورد مقایسه در بخش ۲ ارائه شده است. برای ارزیابی کارایی، از معیارهای رایج طبقه بندی شامل دقت^{۱۶} و فراخوانی^{۱۷} استفاده شده است. دقت نشان می دهد چند درصد از نظراتی که توسط الگوریتم یادگیری ماشین برچسب مثبت (یا منفی) خورده اند، در واقع برچسب مثبت (یا منفی) داشته اند. فراخوانی بیانگر درصد نظرات مثبت (یا منفی) است که به درستی توسط الگوریتم برچسب زده شده اند.

جدول (۹): مقایسه کارایی روش های پیشنهادی با روش های موجود

روش	دقت	فراخوانی
معماری پیشنهادی اول	۸۴/۵۱	۸۴/۰۵
معماری پیشنهادی دوم	۸۸/۵۵	۸۷/۴۸
1C-CNN [23]	۸۲/۹	۷۹/۴۵
1C-CNN [24]	۸۵/۳۳	۸۴/۴۵
Bi-LSTM [25]	۸۶/۸۷	۸۶/۰۳
Voting-Ensemble [26]	۸۲/۳۴	۸۱/۰۳

همان طور که در جدول (۹) مشاهده می شود، معیارهای دقت و فراخوانی معماری پیشنهادی اول در مقایسه با معماری مشابه [۲۳] که در آن شبکه عصبی پیچشی تنها یک کانال دارد، به ترتیب به میزان ۱/۶ و ۴/۶ درصد بهبود داشته است. به طور مشابه، درصد بهبود دقت و فراخوانی برای معماری پیشنهادی دوم در مقایسه با معماری پایه مشابه آن [۲۴]، ۳/۲ و ۳ درصد است. این بهبود نشان می دهد که ویژگی های مستخرج از واژه نامه حسی کمک شایانی

¹⁷ Recall¹⁶ Precision

۵. نتیجه گیری

در این مقاله ابتدا به گسترش واژه‌نامه‌های حسی موجود در زبان فارسی با عبارات محاوره‌ای و غیررسمی رایج در متون نوشته شده در شبکه‌های اجتماعی پرداخته شد. در این راستا دو روش برای تعیین قطبیت عبارات محاوره‌ای پیشنهاد شد که از قابلیت‌های ChatGPT در درک معنا و قطبیت عبارات محاوره‌ای فارسی بهره می‌برند. حاصل هر یک از این دو روش، مجموعه‌ای از عبارات محاوره‌ای در زبان فارسی است که برچسب قطبیت دارند. سپس از یک شبکه عصبی پیچشی دوکاناله که ورودی کانال‌های آن متن نظر و دنباله برچسب قطبیت کلمات متن است برای طبقه‌بندی نظرات در دو گروه مثبت و منفی استفاده شد. نتایج ارزیابی‌های انجام گرفته نشان می‌دهد که با افزودن عبارات محاوره‌ای به واژه‌نامه‌های حسی موجود، صحت الگوریتم تحلیل احساسات مبتنی بر واژه‌نامه افزایش می‌یابد. همچنین ترکیب روش مبتنی بر واژه‌نامه با یادگیری عمیق منجر به بهبود کارایی طبقه‌بندی احساسات می‌شود.

بر اساس تحلیل نتایج حاصل از روش پیشنهادی، عوامل اصلی بروز خطا در مدل طبقه‌بندی احساسات، به صورت زیر دسته‌بندی شده‌اند. این موارد می‌توانند به عنوان محورهای اصلی در تحقیقات آتی مورد توجه قرار گیرند: (۱) عدم پوشش کامل عبارات محاوره‌ای رایج در واژه‌نامه حسی، (۲) وجود تعداد قابل توجهی از غلط‌های املایی و استفاده از نگارش عامیانه در نظرات، (۳) وجود متونی با تعداد بالای کلمات مثبت (یا منفی) که در عین حال، قطبیت کلی آن‌ها منفی (یا مثبت) است، (۴) اشتباهات ناشی از عملکرد ChatGPT در تعیین قطبیت برخی واژه‌ها و (۵) خطاهای ناشی از ابزارهای پردازش زبان طبیعی نظیر تجزیه‌گر نحوی و ریشه‌یاب.

یکی دیگر از چالش‌هایی که روش‌های پیشنهادی با آن مواجه شدند، تعیین قطبیت واژه‌هایی مانند "سبک" و "رند" است که

به آموزش شبکه کرده است. همچنین معماری پیشنهادی دوم از روش مبتنی بر Bi-LSTM [۲۵] که تنها متکی بر ویژگی‌های متنی است و برای جاسازی کلمات از پیکره‌های متنی رسمی بهره می‌برد عملکرد بهتری داشته است. برای نمونه در مقایسه با مدل bi-LSTM دقت و فراخوانی روش پیشنهادی به ترتیب ۱/۶۸ و ۱/۴۵ درصد بهبود داشته است. همچنین معماری پیشنهادی دوم در مقایسه با روش Voting-Ensemble [۲۶] که مبتنی بر تجميع نتایج چهار مدل یادگیری عمیق ساده است، حدود ۶ درصد دقت و فراخوانی طبقه‌بندی احساسات را بهبود داده است.

در آزمایش پایانی، کارایی روش‌های پیشنهادی بر روی یک مجموعه داده دیگر به نام DeepSentiPers [۲۵] که حاوی نظرات کاربران دیجی‌کالا می‌باشد ارزیابی شده است. نظرات این مجموعه داده در مقایسه با مجموعه داده توثیق‌های فارسی نگارش رسمی‌تر و اندازه بزرگتری (متوسط طول نظرات در این مجموعه داده ۲۳/۷۲ کلمه است) دارند. برای این منظور، بخشی از این مجموعه داده شامل ۱۸۵۴ نظر انتخاب شده است که ۲۰ درصد آن به صورت تصادفی جدا شده و به عنوان مجموعه تست در نظر گرفته شده است. همان‌طور که در جدول (۱۰) مشاهده می‌شود، معماری پیشنهادی دوم در مقایسه با سایر روش‌ها به کارایی بالاتری در تحلیل احساسات این مجموعه داده دست یافته است.

جدول (۱۰): مقایسه کارایی روش‌ها بر روی مجموعه داده

DeepSentiPers		
روش	دقت	فراخوانی
معماری پیشنهادی اول	۶۰/۹۸	۶۲/۲۶
معماری پیشنهادی دوم	۶۴/۳	۶۵/۵
1C-CNN [23]	۵۶/۲۷	۶۰/۱۱
1C- CNN [24]	۵۷/۰۹	۶۱/۱۹
Bi-LSTM [25]	۵۸/۶۷	۶۱/۹۹
Voting-Ensemble [26]	۵۶/۸۳	۶۴/۴۲

- [۷] م. رسولی و و. کیانی، "مروری بر روش‌های یادگیری عمیق برای طبقه‌بندی هیجان در متن: پیشرفت‌ها، چالش‌ها و فرصت‌ها"، محاسبات نرم، ۱۴۰۲، doi.org/10.22052/scj.2023.248812.1126
- [8] E. Asgarian, M. Kahani, and Sh Sharifi, "HesNegar: Persian Sentiment WordNet," *Signal and Data Processing*, vol. 15, no. 1, pp. 71-86, 2018, doi.org/10.29252/jsdp.15.1.71 [In Persian].
- [9] B. Sabeti, P. Hosseini, G. Ghassem-Sani, and S. A. Mirroshandel, "LexiPers: An ontology based sentiment lexicon for Persian," *arXiv preprint arXiv:1911.05263*, 2019, doi.org/10.48550/arXiv.1911.05263.
- [10] F. Amiri, S. Scerri, and M. Khodashahi, "Lexicon - based sentiment analysis for Persian Text," in *Proceedings of the International Conference Recent Advances in Natural Language Processing*, 2015, pp. 9-16, doi.org/10.13140/RG.2.1.2537.8327.
- [11] K. Dashtipour, A. Raza, A. Gelbukh, R. Zhang, E. Cambria, and A. Hussain, "PerSent 2.0: Persian Sentiment Lexicon Enriched with Domain-Specific Words", in *Advances in Brain Inspired Cognitive Systems: 10th International Conference*, Guangzhou, China, 2019, pp. 497-509, doi.org/10.1007/978-3-030-39431-8_48.
- [12] K. Dashtipour, M. Gogate, A. Gelbukh, and A. Hussain, "Extending persian sentiment lexicon with idiomatic expressions for sentiment analysis," *Social Network Analysis and Mining*, vol. 12, pp. 1-13, 2022, doi.org/10.1007/s13278-021-00840-1.
- [13] S. Baccianella, A. Esuli, and F. Sebastiani, "Sentiwordnet 3.0: An enhanced lexical resource for sentiment analysis and opinion mining," in *Lrec*, vol. 10, 2010, pp. 2200-2204.
- [14] E. Cambria, Y. Li, F. Z. Xing, S. Poria, and K. Kwok, "SenticNet 6: Ensemble application of symbolic and subsymbolic AI for sentiment analysis," in *Proceedings of the 29th ACM international conference on information & knowledge management*, 2020, pp. 105-114, doi.org/10.1145/3340531.3412003.
- [15] I. Dehdarbehbahani, A. Shakery, and H. Faili, "Semi-supervised word polarity identification in resource-lean languages," *Neural networks*, vol. 58, pp. 50-59, 2014, doi.org/10.1016/j.neunet.2014.05.018.
- [16] M. Darwich, S. A. Mohd, N. Omar, and N. A. Osman, "Corpus-Based Techniques for Sentiment Lexicon
- قطبیت آنها با توجه به متن اطراف می‌تواند متفاوت باشد. به عنوان کار آتی پیشنهاد می‌شود از قابلیت‌های ChatGPT برای تعیین قطبیت واژه‌های مبهم استفاده شود. همچنین بکارگیری ویژگی‌هایی مانند نقش نحوی و معنایی کلمات جهت بهبود دقت شبکه‌های عصبی چندکاناله برای تحلیل احساسات پیشنهاد می‌شود.
- ### مراجع
- [۱] ا. خلف بیگی، ع. بشیری موسوی و س. قارلقی، "بکارگیری مدل تحلیل احساسات در سطح حروف مبتنی بر شبکه عصبی روی نظرات فارسی ثبت شده در شبکه‌های اجتماعی و فروشگاه‌های اینترنتی"، محاسبات نرم، جلد ۱۱، شماره ۲، صفحات ۱۱۸-۱۳۳، ۱۴۰۱، doi.org/10.22052/scj.2023.248311.1094
- [۲] ک. جهان بین و م. ع. زارع چاهوکی، "تحلیل احساسات رمازرها با یادگیری انتقالی شات صفر"، محاسبات نرم، ۱۴۰۳، doi.org/10.22052/scj.2025.255169.1258
- [3] Z. Rajabi and M. Valavi, "A survey on sentiment analysis in Persian: a comprehensive system perspective covering challenges and advances in resources and methods," *Cognitive Computation*, vol. 13, no. 4, pp. 882-902, 2021, doi.org/10.1007/s12559-021-09886-x.
- [4] Z. Ayeste and S. Noferesti, "A semantic approach based on domain knowledge for polarity shift detection using distant supervision," *Progress in Artificial Intelligence*, vol. 11, no. 2, pp. 169-180, 2022, doi.org/10.1007/s13748-021-00267-x.
- [5] E. Asgarian, M. Kahani, and S. Sharifi, "The impact of sentiment features on the sentiment polarity classification in Persian reviews," *Cognitive Computation*, vol. 10, no. 1, pp. 117 - 135, 2018, doi.org/10.1007/s12559-017-9513-1.
- [6] K. Dashtipour, A. Hussain, Q. Zhou, A. Gelbukh, A. Y. Hawalah, and E. Cambria, "PerSent: A freely available Persian sentiment lexicon," in *Advances in Brain Inspired Cognitive Systems*, Beijing, China, 2016, pp. 310-320, doi.org/10.1007/978-3-319-49685-6_28.

- sentiment corpus,” *arXiv preprint arXiv:2004.05328*, 2020, doi.org/10.48550/arXiv.2004.05328.
- [26] S. Eyvazi-Abdoljabbar, S. Kim, M. R. Feizi-Derakhshi, Z. Farhadi, and D. A. Mohammed, “An ensemble-based model for sentiment analysis of persian comments on instagram using deep learning algorithms,” *IEEE Access*, vol. 12, pp. 151223-151235, 2024, doi.org/10.1109/ACCESS.2024.3473617.
- [27] O. Davar, G. Dar, and F. Ghasemian, “DeepSentiParsBERT: a deep learning model for Persian sentiment analysis using ParsBERT,” in *28th International Computer Conference, Computer Society of Iran (CSICC)*, 2023, pp. 1-5, doi.org/10.1109/CSICC58665.2023.10105414.
- [28] M. Farahani, M. Gharachorloo, M. Farahani, and M. Manthouri, “Parsbert: Transformer-based model for persian language understanding,” *Neural Processing Letters*, vol. 53, pp. 3831-3847, 2021, doi.org/10.1007/s11063-021-10528-4.
- [29] F. Ariai, M. T. Mahmoudi, and A. Moeini, “Enhancing aspect-based sentiment analysis with ParsBERT in Persian language,” *arXiv preprint arXiv:2502.01091*, 2025, doi.org/10.48550/arXiv.2502.01091.
- [30] M. Masumi, S. S. Majd, M. Shamsfard, and H. Beigy, “Fabert: Pre-training bert on persian blogs,” *arXiv preprint arXiv:2402.06617*, 2024, https://doi.org/10.48550/arXiv.2402.06617.
- [31] A. Shokri (2024. October. 1). Available: <https://github.com/semnan-university-ai/persian-slang>.
- [32] Dataheart (2024. August. 10). Available: <http://dataheart.ir/category/67/> نظر کاوی.
- Generation: A Review,” *Journal of Digital Information Management*, vol. 17, no. 5, pp. 296, 2019, doi.org/10.6025/jdim/2019/17/5/296-305.
- [17] R. Dehkharghani, “Sentifars: A persian polarity lexicon for sentiment analysis,” *ACM Transactions on Asian and Low-Resource Language Information Processing (TALLIP)*, vol. 19, no. 2, pp. 1-12, 2019, doi.org/10.1145/3345627.
- [18] K. Dashtipour, M. Gogate, A. Adeel, H. Larijani, and A. Hussain, “Sentiment analysis of persian movie reviews using deep learning,” *Entropy*, vol. 23, no. 5, pp. 596, 2021, doi.org/10.3390/e23050596.
- [19] E. Grave, P. Bojanowski, P. Gupta, A. Joulin, and T. Mikolov, “Learning Word Vectors for 157 Languages,” *arXiv preprint arXiv:1802.06893*, 2018, doi.org/10.48550/arXiv.1802.06893.
- [20] M. B. Dastgheib, S. Koleini, and F. Rasti, “The application of deep learning in persian documents sentiment analysis,” *International Journal of Information Science and Management (IJISM)*, vol. 18, no. 1, pp. 1-15, 2020.
- [21] O. Davar, G. Dar, and F. Ghasemian, “DeepSentiParsBERT: a deep learning model for Persian sentiment analysis using ParsBERT,” in *28th International Computer Conference, Computer Society of Iran (CSICC)*, 2023, pp. 1-5, doi.org/10.1109/CSICC58665.2023.10105414.
- [22] M. Dehghani, and Z. Yazdanparast, “Sentiment analysis of Persian political tweets using ParsBERT embedding model with convolutional neural network,” in *9th International Conference on Web Research (ICWR)*, 2023, pp. 20-25, doi.org/10.1109/ICWR57742.2023.10139063.
- [۲۳] م. روحانیان، م. صالحی، ع. درزی و و. رنجبر، “تحلیل احساس در رسانه‌های اجتماعی فارسی با رویکرد شبکه عصبی پیچشی،” *مهندسی برق و مهندسی کامپیوتر ایران*، جلد ۱۸، شماره ۱، صفحات ۵۹-۶۶، ۱۳۹۹.
- [24] M. Vazan, and J. Razmara, “Jointly modeling aspect and polarity for aspect-based sentiment analysis in persian reviews,” *arXiv preprint arXiv:2109.07680*, 2021, doi.org/10.48550/arXiv.2109.07680.
- [25] J. P. R. Sharami, P. A. Sarabestani, and S. A. Mirroshandel, “Deepsentipers: Novel deep learning models trained over proposed augmented persian