

چارچوبی برای تحلیل احساسات داده‌های متنی و شکلک‌ها در شبکه‌های اجتماعی

مرتضی شریفی^۱، کارشناسی ارشد، حجت اله حمیدی^۲، دانشیار

۱- گروه فناوری اطلاعات- دانشگاه صنعتی خواجه نصیرالدین طوسی - تهران - ایران - m.sharifi@email.kntu.ac.ir

۲- گروه فناوری اطلاعات- دانشگاه صنعتی خواجه نصیرالدین طوسی - تهران - ایران - h_hamidi@kntu.ac.ir

چکیده: امروزه افراد به‌منظور افزایش احساس خود در متن یا خلاصه کردن عبارات، از شکلک‌ها استفاده می‌کنند، تکنیک‌های قبلی یادگیری ماشین فقط شامل طبقه‌بندی متن، شکلک یا تصاویر، به‌تنهایی بوده است، جایی که شکلک‌ها با متن اکثر اوقات نادیده گرفته شده است. تحقیقات نشان می‌دهد یادگیری عمیق به‌ندرت در تحلیل احساسات بر روی ترکیب داده‌های متن و شکلک اعمال شده است. بنابراین، این مقاله متن و شکلک‌ها را به‌طور جداگانه و به‌صورت ترکیبی برای یافتن احساسات، تجزیه و تحلیل کرده است. بدین‌صورت که ابتدا یک تجزیه و تحلیل مقایسه‌ای از شبکه‌های عصبی کانولوشنی و شبکه‌های حافظه طولانی کوتاه‌مدت انجام می‌گردد، که جهت افزایش دقت تعبیه کلمات، بردارهای تعبیه کلمات از قبل آموزش‌دیده word2vec, glove, BERT و RoBERTa مورد استفاده قرار گرفته است. سپس یک مدل ترکیبی مبتنی بر توجه ارائه می‌گردد که از شبکه‌های عصبی کانولوشنی و شبکه‌های حافظه طولانی کوتاه‌مدت دوطرفه استفاده می‌کند. نتایج نشان می‌دهد که در نظر گرفتن شکلک‌ها به عنوان ویژگی، سبب افزایش دقت مدل شده است، همچنین مدل پیشنهادی بر اساس چهار معیار ارزیابی دقت، صحت، فراخوانی و امتیاز f1 از مدل‌های پایه و کارهای قبلی بررسی شده، بهتر عمل کرده، که دقت مدل پیشنهادی بر روی مجموعه داده بدون شکلک ۸۵/۸۰ بوده و برای مجموعه داده شامل شکلک ۹۰/۱۸ است.

واژه‌های کلیدی: تجزیه و تحلیل احساسات، عقیده کاوی، پردازش زبان طبیعی، یادگیری عمیق.

A framework for sentiment analysis of textual data and emoticons on social networks

Morteza Sharifi, MSc Student¹, Hojatollah Hamidi, Associate Professor²

1- Information Technology Engineering Group, K.N.Toosi University of Technology, Tehran, Iran, Email: m.sharifi@email.kntu.ac.ir

2- Information Technology Engineering Group, K.N.Toosi University of Technology, Tehran, Iran, Email: h_hamidi@kntu.ac.ir

Abstract: Nowadays, people use emoticons in the text to increase their feelings or summarize expressions. Earlier machine learning techniques only involve the classification of text, emoticons, or images solely, whereas emoticons with text have most of the time been neglected. Research shows deep learning was rarely applied in sentiment analysis on text and emoticon data combination. Therefore, this article has analyzed the text and emoticons separately and in combination to find the emotions. First a comparative analysis of convolutional neural networks and long-short term memory networks is performed, to increase the accuracy of word embedding, pre-trained word embedding vectors word2vec, glove, BERT and RoBERTa were used. Then, a new hybrid attention-based model is presented, which uses convolutional neural networks and Bidirectional long-short term memory networks. The results show that considering emoticons as a feature increases the accuracy of the model, also the proposed model has performed better than the basic models and previously reviewed works based on the four evaluation criteria of accuracy, precision, recall, and f1 score that accuracy of the proposed model on the dataset without emoticons is 85.80 and on the dataset including emoticons is 90.18.

Keywords: Sentiment analysis, Opinion mining, Natural language processing, Deep learning.

تاریخ ارسال مقاله: ۱۴۰۱/۰۶/۱۷

تاریخ اصلاح مقاله: ۱۴۰۲/۰۲/۱۶

تاریخ پذیرش مقاله: ۱۴۰۲/۰۴/۲۱

نام نویسنده مسئول: حجت اله حمیدی

نشانی نویسنده مسئول: ایران - تهران - دانشگاه صنعتی خواجه نصیرالدین طوسی - گروه آموزشی فناوری اطلاعات.

۱- مقدمه

تجزیه و تحلیل احساسات^۱ در نظارت بر شبکه‌های اجتماعی بسیار ارزشمند است زیرا به شما امکان می‌دهد بینشی در موضوعات خاص بدست آورید [۱]. رسانه‌های اجتماعی مانند توئیتر، فیس‌بوک و یوتیوب در حال حاضر از محبوب‌ترین مکان‌هایی هستند که مشتریان می‌توانند محصولات و خدمات جدید را در بازارهای مختلف مورد بحث و بررسی قرار دهند [۲]. به‌طور سنتی، تجزیه و تحلیل احساسات بر روی متن انجام شده است، اما مقدار زیادی از داده‌ها اکنون به‌عنوان نظرات، شامل شکلک‌ها هم هستند. با بررسی این داده‌ها، می‌توان احساسات عمومی نسبت به یک موضوع خاص را تحلیل، بررسی و کشف کرد [۳]. مدل تجزیه و تحلیل احساسات، بسته به اینکه شکلک‌ها^۲ را در نظر می‌گیرد یا آن‌ها را نادیده می‌گیرد، حالت‌های مختلفی را برای هر توئیت، تولید می‌کند [۴]. در سال‌های اخیر، نشان داده‌شده است که استفاده از شبکه‌های عصبی عمیق در طبقه‌بندی متن مؤثر بوده است [۵]، همچنین تحقیقات اخیر در مورد کشف قطبیت شامل تکنیک‌های شبکه عصبی کانولوشنی^۳، حافظه طولانی کوتاه‌مدت^۴ و جاسازی کلمات^۵ است [۶]. شبکه عصبی کانولوشنی در ابتدا برای حوزه بینایی ماشین ساخته شد. پس‌از آن، شبکه‌های عصبی کانولوشنی برای اهداف پردازش زبان طبیعی مورد استفاده قرار گرفتند و برای طبقه‌بندی متن و تحلیل احساسات به نتایج عالی دست‌یافتند [۷، ۸]. حافظه طولانی کوتاه‌مدت، یک مدل یادگیری عمیق مؤثر در تجزیه و تحلیل دنباله‌های طولانی، برای طبقه‌بندی احساسات در زمینه پردازش زبان طبیعی مورد استفاده قرار گرفته است [۹]. روش‌های یادگیری عمیق از جاسازی کلمات برای یادگیری روابط معنایی کلمات به‌منظور بهبود عملکرد مدل استفاده می‌کنند [۱۰]. در این مقاله اهمیت قرار دادن شکلک‌ها در پیام‌های توئیتر بررسی می‌شود و یک مدل ترکیبی مبتنی بر توجه شبکه عصبی کانولوشنی و شبکه حافظه طولانی کوتاه‌مدت پیاده‌سازی می‌گردد.

تحلیل احساسات زیرمجموعه پردازش زبان طبیعی^۶ است. تحقیقات بسیار کمی در مورد شکلک‌ها برای یافتن احساسات انجام شده است [۱۱]. زمینه برای گسترش بیشتر در حوزه تحلیل احساسات با متن و شکلک وجود دارد. اهداف این تحقیق عبارتند از:

- ۱) بررسی تأثیر شکلک‌ها در تجزیه و تحلیل احساسات داده‌های شبکه‌های اجتماعی
 - ۲) بررسی دقت طبقه‌بندی تجزیه و تحلیل احساسات با استفاده از الگوریتم‌های یادگیری عمیق
 - ۳) بررسی تأثیر بردارهای تعبیه کلمات از قبل آموزش‌دیده مختلف در افزایش دقت مدل
 - ۴) ارائه یک مدل ترکیبی مبتنی بر توجه با استفاده از الگوریتم‌های یادگیری عمیق برای تجزیه و تحلیل احساسات
- در ادامه، مقاله به شکل زیر سازماندهی شده است: در بخش دوم، مروری بر تحقیقات پیشین بیان شده است، در بخش سوم، سطوح

تجزیه و تحلیل احساسات بررسی شده، در بخش چهارم، پیاده‌سازی روش پیشنهادی ارائه شده است، در بخش پنجم، نتایج مدل‌ها، در بخش ششم، ارزیابی نتایج و در بخش پایانی، نتیجه‌گیری و پیشنهاد برای تحقیقات آتی، بیان شده است.

۲- مروری بر تحقیقات پیشین

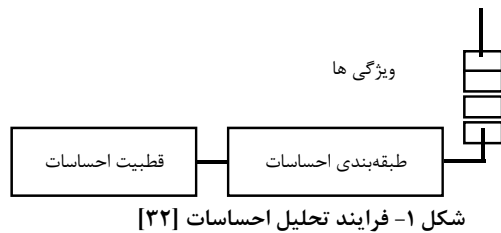
فاکتور تازگی، همراه با کاربردهای گسترده و در دسترس بودن داده‌های رسانه‌های اجتماعی محققان زیادی را به حوزه تجزیه و تحلیل احساسات جذب کرده است. تحقیقات جدید تجزیه و تحلیل احساسات اکنون در حوزه‌ها و زبان‌های مختلفی انجام شده است [۱۲-۱۳]. وان و گائو از یک روش جمعی طبقه‌بندی احساسات با کمک روش‌های مختلف طبقه‌بندی مانند بیز ساده، ماشین بردار پشتیبان، درخت تصمیم‌گیری و جنگل تصادفی استفاده کردند، آن‌ها شش تکنیک طبقه‌بندی را آزمایش کردند و با توجه به رویکرد جمعی آموزش داده‌شده و تست شدند [۱۴]. پاک و پروبک داده‌های میکرو بلاگ را طبقه‌بندی کرده و احساسات مثبت، منفی و خنثی را تعیین کردند. سهم اصلی جمع‌آوری خودکار مجموعه داده‌ها و تحلیل احساسات است، تکنیک پیشنهادی آن‌ها کارآمد بوده و عملکرد بهتری نسبت به روش‌های پیشنهادی قبلی دارد [۱۵]. ویزلاو و لونی داده‌های توئیتر را با استفاده از روش‌های یادگیری ماشین طبقه‌بندی کرد. تعداد محدودی از شکلک‌های استاندارد از داده‌های توئیتر شناسایی شد و احساسات مربوطه را بدست آورد [۱۶]. کولومبیس و همکاران، ۲۰۸۰ توئیت به زبان هلندی را که حاوی شکلک هستند ارزیابی کردند. نشان دادند شرح سطح پاراگراف برای احساسات دلالت شده توسط شکلک‌ها به طرز قابل توجهی از دقت طبقه‌بندی احساسات پیشی می‌گیرد که نشان می‌دهد مشارکت شکلک‌ها بر احساسات منتقل شده توسط نشانه‌های متنی غلبه می‌کند [۱۷]. تعبیه کلمه به‌عنوان نمایش کلمات جهت تحلیل متن، معمولاً در قالب یک بردار که معنای کلمه را رمزگذاری می‌کند، تعریف می‌شود. به این ترتیب انتظار می‌رود کلماتی که در یک فضای برداری نزدیک‌تر ظاهر می‌شوند معانی مشابهی داشته باشند. مدل‌هایی مانند word2vec و Glove معنای لغوی کلمات را در نظر می‌گیرند [۱۸]. موقعیت‌هایی وجود دارد که در آن‌ها جنبه ترتیبی متن مهم جلوه می‌کند. از این نظر، جاسازی‌های کلمه، تکنیک‌های تفسیر متنی هستند که تلاش می‌کنند زمینه^۷ و موضوع کلمات را در قالب بردارهایی با طول ثابت نشان دهند [۱۹]. Word2vec یکی از پرکاربردترین الگوریتم‌های جاسازی کلمه است که نتایج پیشرفته‌ای را در کاربردهای متعدد پردازش زبان طبیعی ارائه کرده است [۲۰]. لو و همکاران از Word2Vec برای نمایش کلمات مجموعه داده به‌عنوان ورودی خودرمزگذار بازگشتی جهت تحلیل احساسات در سطح سند^۸ استفاده کردند. نتایج نشان داد که مدل پیشنهادی نسبت به برخی از روش‌ها با نظرات مانند ماشین بردار پشتیبان^۹ و بیز ساده^{۱۰} عملکرد بهتری دارد [۲۱]. GloVe یک مدل رگرسیون لگاریتمی دوسویه است که تلاش می‌کند با ترکیب مزایای روش‌های مبتنی بر شمارش و مبتنی بر پیش‌بینی، بر محدودیت‌های روش‌های وان-هات^{۱۱}، فروانی کلمه-معکوس فراوانی سند^{۱۲} و Word2Vec غلبه کند [۲۲]. شارما و همکاران مدلی را با GloVe و شبکه

ویژگی‌های مهم را در هر گروه افزایش می‌دهد در حالی که موارد کم‌اهمیت‌تر را تضعیف می‌کند. نتایج تجربی نشان می‌دهد که مکانیسم ارتقای گروهی می‌تواند نتایج امیدوارکننده‌ای را در مورد تخصیص مقادیر وزنی مناسب به ویژگی‌های زمینه به همراه داشته باشد و معماری پیشنهادی می‌تواند برای تحلیل احساسات بهتر عمل کند [۲۸]. محمد امان‌الله و همکاران یک الگوریتم و روش برای تجزیه و تحلیل احساسات با استفاده از متن و شکلک ارائه دادند. در این کار، هر دو حالت داده به صورت ترکیبی و جداگانه با استفاده از الگوریتم‌های یادگیری ماشین و یادگیری عمیق تحلیل کردند. نشان دادند که در تحلیل احساسات الگوریتم‌های یادگیری عمیق، شبکه‌های عصبی کانولوشنی و شبکه‌های حافظه طولانی کوتاه‌مدت بهتر از الگوریتم‌های یادگیری ماشین عمل می‌کنند [۲۹]. بنابراین، در این تحقیق متن توپیت از مجموعه داده Sentiment140، [۲۵-۲۸]، با استفاده از تکنیک‌های یادگیری عمیق شبکه‌های عصبی کانولوشنی و شبکه‌های حافظه طولانی کوتاه‌مدت به همراه بردارهای تعبیه کلمات از قبل آموزش دیده در نظر گرفته شد تا بررسی گردد که آیا استفاده از شکلک‌ها می‌تواند دقت پیش‌بینی را بهبود دهد، در نهایت یک مدل ترکیبی مبتنی بر توجه شبکه‌های عصبی کانولوشنی با شبکه‌های حافظه طولانی کوتاه‌مدت پیشنهاد می‌گردد که عملکرد بهتری را نسبت به کارهای بررسی شده ارائه می‌کند. جدول ۱، تعدادی از مطالعات مهم در زمینه تحلیل احساسات را بیان می‌کند.

جدول ۱: تعدادی از مطالعات مهم در زمینه تجزیه و تحلیل احساسات

نویسنده و سال	الگوریتم و روش / محیط کار	هدف و یافته‌ها
ویسلو و ولنی (۲۰۱۶)	طبقه‌بندی احساسات سطح سند، تجزیه و تحلیل نمادها/ twitter	تجزیه و تحلیل نمادها که از شکلک‌ها و کاراکترهای ایموجی استفاده می‌کند، می‌تواند به طور قابل توجهی دقت در شناخت انواع احساسات را افزایش دهد.
محمد احسان بصیری و همکاران (۲۰۲۰)	CNN BI-LSTM BI-GRU ATTENTION /MECHANISM Sentiment140 Airline dataset	ارائه یک مدل ترکیبی یادگیری عمیق برای تحلیل احساسات و نتایج نشان می‌دهد که روش پیشنهادی روی دیتاست‌های شامل توپیت‌های طولانی‌تر بهتر از دیتاست با توپیت‌های کوتاه عمل می‌کند.
محمد امان‌الله و همکاران (۲۰۲۰)	IDF-TF SVM NB FF CNN /LSTM twitter	این تحقیق نشان می‌دهد که هر زمان شکلک استفاده می‌شود، احساسات مرتبط با آن بر احساس منتقل شده توسط تجزیه و تحلیل داده‌های متنی غلبه می‌کند. همچنین، الگوریتم‌های یادگیری عمیق بهتر از الگوریتم‌های یادگیری ماشین هستند.
مرجان کامیاب و همکاران (۲۰۲۱)	TF-IDF CNN LSTM ATTENTION /MECHANISM US-airline-Sentiment140-MV SA4A	ارائه یک مدل مبتنی بر توجه جدید که از CNN با LSTM استفاده می‌کند. دقت مدل پیشنهادی نسبت به روش‌های پایه بهبود یافته است.
ماریا موزنو و همکاران (۲۰۲۱)	SVM CNN LSTM -BERT Sentiment140 Tweets Airline	نتایج نشان می‌دهد که قابلیت اطمینان مدل‌های ترکیبی در بین تمام مدل‌های آزمایش شده برای تحلیل قطبیت احساسات بهتر عمل می‌کند.

عصبی بازگشتی جهت تحلیل احساسات ارائه کرده‌اند، که این مدل علاوه بر طبقه‌بندی دودویی برای طبقه‌بندی چند کلاسه هم می‌تواند مورد استفاده قرار بگیرد [۲۳]. ورشنی و همکاران در [۲۴]، یک تحلیل مقایسه‌ای از شبکه‌های LSTM، حافظه طولانی کوتاه-مدت دوطرفه^{۱۳} و نمایش رمزگذار دوطرفه از ترانسفورمرها^{۱۴} بر روی دو مجموعه داده انجام دادند. نتایج نشان داد که BERT از نظر دقت، صحت، بازخوانی و امتیاز f1 از الگوریتم‌های دیگر عملکرد بهتری دارد. در این مطالعه یک مدل ترکیبی برای تجزیه و تحلیل احساسات پیشنهاد می‌گردد که دقت را افزایش می‌دهد. روش‌های زیادی برای ساخت مدل‌های ترکیبی وجود دارد، برای مثال دنگ و همکاران در [۲۵] عملکرد ترکیب ماشین بردار پشتیبان، CNN و LSTM را با استفاده از تکنیک‌های جاسازی، Word2vec و BERT، روی هشت مجموعه داده متنی از توپیت‌ها آزمایش کردند. پس از آن، مدل ترکیبی تولید شده را با مدل‌های تک مقایسه کردند. نتایج نشان داد که قابلیت اطمینان مدل‌های ترکیبی در بین تمام مدل‌های آزمایش شده برای تحلیل قطبیت احساسات بهتر عمل می‌کند و ترکیب مدل‌های یادگیری عمیق با تکنیک SVM نتایج بهتری نسبت به استفاده از یک مدل فردی برای انجام تحلیل احساسات به همراه دارد. در مطالعه‌ای دیگر نویسندگان در [۲۶] یک مدل عمیق شبکه عصبی بازگشتی با شبکه عصبی کانولوشنی دوجهته مبتنی بر توجه برای تحلیل احساسات در سطح سند پیشنهاد می‌کنند. مدل پیشنهادی با استفاده از دولا به LSTM و GRU دوطرفه مستقل، هر دو زمینه گذشته و آینده را با در نظر گرفتن جریان اطلاعات زمانی در هر دو جهت استخراج می‌کند. همچنین مکانیسم توجه بر روی خروجی لایه‌های دوطرفه برای تأکید بیشتر یا کمتر بر کلمات مختلف اعمال می‌شود. نتایج نشان می‌دهد مدل پیشنهادی نسبت به کارهای بررسی شده عملکرد بهتری دارد و به نتایج پیشرفته‌ای در طبقه‌بندی توپیت‌های طولانی و کوتاه دست می‌یابد. نویسندگان در [۲۷] یک مدل مبتنی بر توجه را با استفاده از شبکه‌های عصبی کانولوشنی و حافظه طولانی کوتاه‌مدت پیشنهاد می‌کنند که از وزن دهی TF-IDF و رویکرد از پیش-آموزش دیده GloVe برای استخراج اطلاعات معنی‌دار استفاده می‌کند. علاوه بر این، یک مکانیسم توجه در انتهای لایه‌های CNN برای جلب توجه بیشتر یا کمتر به کلمات مختلف به کار گرفته شده است. کارایی مدل بر روی چهار مجموعه داده استاندارد انگلیسی، با ماتریس‌های عملکرد مختلف ارزیابی شده و کارایی با مدل‌ها و رویکردهای پایه موجود مقایسه شده است. نتایج آزمایش نشان می‌دهد که روش پیشنهادی به طور قابل توجهی بهتر از مدل‌های پیشرفته عمل می‌کند. آیتوگ اونان یک معماری شبکه عصبی بازگشتی کانولوشنی دوطرفه را پیشنهاد می‌کند که از دولا به دوطرفه LSTM و GRU استفاده می‌کند تا با اتصال دولا به پنهان از جهت‌های مخالف به یک زمینه، هم زمینه‌های گذشته و هم آینده را استخراج کند. مکانیسم ارتقای گروهی بر روی ویژگی‌های استخراج شده توسط لایه‌های دو جهته به کار گرفته شده است، که ویژگی‌ها را به کلاس‌های متعدد تقسیم می‌کند و اهمیت



۲-۳- تکنیک های تحلیل احساسات

روش های تجزیه و تحلیل احساسات مبتنی بر یادگیری ماشین، مبتنی بر واژه نامه و ترکیبی است. در روش یادگیری ماشین از مجموعه داده ویژگی ها را استخراج می کنیم و این ویژگی ها به طبقه بندی قطبیت جمله ورودی ناشناخته کمک می کند. روش های یادگیری ماشین به یادگیری تحت نظارت و یادگیری بدون نظارت تقسیم می شود [۳۲]. یادگیری با نظارت^{۲۱}: این رویکرد زمانی مورد استفاده قرار می گیرد که داده های برچسب گذاری شده برای آموزش مدل وجود داشته باشد. دو مرحله در یادگیری با نظارت استفاده می شود: اول آموزش مدل و دیگری پیش بینی [۳۳]. در طول آموزش، مجموعه داده با برچسب های آن به الگوریتم طبقه بندی تغذیه می شود که یک مدل را به عنوان خروجی می دهد. پس از آن داده های تست برای پیش بینی طبقه بندی وارد مدل می شوند.

یادگیری بدون نظارت^{۲۲}: این روش زمانی مورد استفاده قرار می گیرد که قابلیت اطمینان داده های برچسب دار دشوار باشد. جمله بر اساس لیست کلمات کلیدی هر دسته طبقه بندی می شود. این الگوریتم ها الگوهای پنهان یا گروه بندی داده ها را بدون نیاز به دخالت انسان کشف می کنند. آتیسا^{۲۳} و همکاران تجزیه و تحلیل احساسات را با استفاده از رویکرد بدون نظارت انجام دادند که در آن توییت ها با استفاده از روش خوشه بندی طیفی به خوشه مثبت و منفی تقسیم شدند. خوشه های طیفی از بیز ساده، ماشین بردار پشتیبان و حداکثر آنتروپی^{۲۴} عملکرد بهتری داشتند [۳۴-۳۷].

یادگیری عمیق برای طبقه بندی احساسات: شبکه عصبی مصنوعی ساختار عصبی مغز انسان را تقلید می کند. واحد اصلی شبکه عصبی نورون است. شامل یک لایه ورودی، یک لایه پنهان و یک لایه خروجی است. شبکه عصبی از تابع خطی $x(i) = A \cdot a(i)$ استفاده می کند. بردار " $a(i)$ " به عنوان ورودی به نورون داده می شود که یک کلمه را در یک سند نشان می دهد. یک وزن " A "، مربوط به هر نورون وجود دارد که برای محاسبه عملکرد مورد استفاده قرار می گیرد. $x(i)$ برای طبقه بندی کلاس استفاده می شود. در شبکه های عصبی مصنوعی، آموزش مدل شامل دو مرحله است: انتشار روبه جلو و انتشار عقبگرد. در انتشار روبه جلو، ورودی به لایه ورودی نورون ها داده می شود که در وزن هایی که اعداد تصادفی هستند ضرب می شود. توابع برای نرمال سازی مقدار خروجی بین ۰ و ۱ استفاده می شوند. سپس خروجی با مقدار مورد نظر مقایسه می شود، اگر بین دو مقدار تفاوت (خطا) وجود داشته باشد، انتشار عقبگرد انجام می شود.

	SemEval	
--	---------	--

۳- سطوح تجزیه و تحلیل احساسات

تحقیقات تجزیه و تحلیل احساسات به طور عمده در سه سطح انجام شده است:

سطح سند، سطح جمله^{۱۵} و سطح جنبه^{۱۶}.

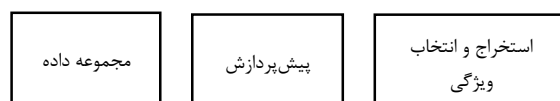
سطح سند: کار در سطح سند این است که طبقه بندی یک سند کامل دارای احساسات منفی است یا مثبت [۳۰]. این سطح از تجزیه و تحلیل به طور ضمنی فرض می کند که هر سند نظرات خود را در مورد یک موجودیت واحد (به عنوان مثال، یک محصول یا خدمت واحد) بیان می کند؛ بنابراین برای اسنادی که چندین نهاد را ارزیابی یا مقایسه می کنند که برای آن ها تجزیه و تحلیل دقیق تر لازم است، قابل استفاده نیست.

سطح جمله: سطح جمله تعیین این است که هر جمله نظر مثبت، منفی یا خنثی را بیان می کند. توجه داشته باشید که «نظر خنثی» معمولاً به معنای «عدم عقیده» است.

سطح جنبه: نه در آنالیزهای سطح سند و نه در آنالیزهای سطح جمله، چیزی که مردم دقیقاً دوست دارند و دوست ندارند، کشف نمی شوند. به عبارت دیگر، آن ها نمی گویند که هر نظر در مورد چیست، یعنی هدف نظر. می توان گفت اگر بتوانیم جمله ای را مثبت طبقه بندی کنیم، همه چیز در جمله می تواند نظر مثبت را بدست آورد. با این حال، این نیز مفید نخواهد بود، زیرا یک جمله می تواند چندین نظر داشته باشد، به عنوان مثال، «پل در این اقتصاد ضعیف بسیار خوب عمل می کند». منطقی نیست که این جمله را مثبت یا منفی طبقه بندی کنیم زیرا در مورد پل مثبت است اما در مورد اقتصاد منفی است. برای به دست آوردن این سطح از نتایج ریز، باید به سطح جنبه برویم. این سطح از تجزیه و تحلیل قبلاً سطح ویژگی نامیده می شد، اکنون تجزیه و تحلیل احساسات مبتنی بر جنبه نامیده می شود [۳۱]. هدف این سطح از تجزیه و تحلیل کشف احساسات موجودیت ها و یا جنبه های آن ها است.

۳-۱- فرایند کلی تجزیه و تحلیل احساسات

تجزیه و تحلیل احساسات شامل پیش پردازش داده ها^{۱۷}، انتخاب ویژگی ها^{۱۸}، طبقه بندی و سپس پیدا کردن قطبیت داده ها است (شکل ۱). پیش پردازش داده ها شامل توکن سازی، حذف کلمات توقف و ریشه یابی^{۱۹} است. توکن سازی عمل شکستن دنباله ای از کلمات به کلمات جداگانه به نام توکن است. کلمات توقف کلماتی هستند (are, am, is, to, in و غیره) که هیچ نظری ندارند، بنابراین حذف آن ها می تواند مفید واقع شود. ریشه یابی، عمل تبدیل اشکال مختلف کلمه، به شکل اصلی و پایه آن مانند helping به help است.



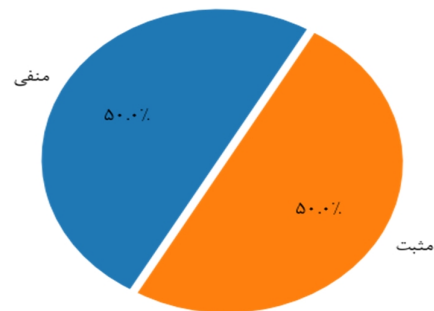
XP	203	:o	648	x3	1540
:3	1627	=]	469	X3	102
xD	1771	0:3	314	=3	48
XD	1830	8)	308	;3	13339
:p	1732	:*	281	:b	81
:O	1606	:]	170	=D	837
:-)	1101	=p	160	:L	667
:]	1019	:o)	83	=/	520
=\	53	:-/	402	:[	211
:@	417	DX	215	;(	343
:('	1411	:/	14646	:	1083
:S	840				

۴- پیاده‌سازی روش پیشنهادی

در این مقاله مدل‌های پایه با استفاده از توییت‌هایی که شکلک‌های آن‌ها حذف شده آموزش دیده و آزمایش شده‌اند. برای هر مدل پایه، یک مدل متناظر آموزش داده شد که شامل شکلک‌ها به عنوان ورودی بود. دقت^{۲۵}، صحت^{۲۶}، فراخوانی^{۲۷} یا پوشش و امتیاز fl مدل‌های پایه مربوطه که از شکلک استفاده نمی‌کنند با مدل‌های شامل شکلک‌ها مقایسه شدند.

۴-۱- مجموعه داده

استفاده از پایگاه‌های داده فارسی، دارای محدودیت‌هایی در مدت زمان پیاده‌سازی و تامین منابع دارد (که به عنوان یکی از کارهای آینده، در ادامه این تحقیق پیشنهاد می‌شود). در این مقاله، برای ارزیابی احساسات در پیام‌های تویتر، از مجموعه داده زبان انگلیسی Sentiment140 استفاده شده است. شکلک‌ها با استفاده از لیستی از شکلک‌های زبان انگلیسی ارائه شده در ویکی‌پدیا تهیه شده است، که در آن هر نماد از شکلک‌ها، با متن مناسب جایگزین شده است. برای مجموعه داده توییت‌ها، در متن و شکلک‌ها، شش ویژگی در نظر گرفته شده است. این شش ویژگی شامل: قطبیت توییت، شناسه توییت، تاریخ توییت، پرس جوی کاربر، متن توییت و کاربری که توییت کرده است، می‌باشد. شکل ۲، توزیع داده‌ها را برای برچسب‌های احساسات، نشان می‌دهد. هر دو برچسب احساسات مثبت و منفی، به‌طور مساوی در این مجموعه داده توزیع شده‌اند.



شکل ۲- توزیع برچسب احساسات مجموعه داده

در این تحقیق، فقط ردیف‌هایی که حداقل دارای یک شکلک باشند در نظر گرفته شده‌اند، علاوه بر این، دیتاست شامل: هشتگ، آدرس‌های اینترنتی، تگ‌های اچ‌تی‌ام‌ال^{۲۸}، کاراکترهای خاص، کاراکترهای خط جدید (n) و کاراکترهای انتهای خط (t) نیز می‌باشد. در جدول ۲، چهل شکلک، به همراه تعداد وقوع آن‌ها در توییت‌ها، نشان داده شده است.

جدول ۲- شکلک‌ها با تعداد وقوع آن‌ها

شکلک	وقوع	شکلک	وقوع	شکلک	وقوع
:)	4790	:('	199	:*	78
xp	5949	:\$	1506	;D	1619

ستونی از متن که شامل شکلک می‌باشد، با استفاده از لیستی از شکلک‌های زبان انگلیسی ذکر شده در ویکی‌پدیا تهیه شده، که در آن هر نماد شکلک با متن مناسب جایگزین شده است. به عنوان مثال، ":]" و "[=" با happy face و "؛)" و sad جایگزین شده‌اند. بیش از ۹۰٪ توییت‌ها یک یا دو شکلک داشتند و در تعداد کمی از توییت‌ها بیش از دو شکلک وجود داشت. تعداد ردیف‌های هر دو مجموعه داده، بر اساس قطبیت مثبت و منفی در جدول ۳ نشان داده شده است.

جدول ۳- تعداد ردیف‌ها در دیتاست پایه و دیتاست شامل شکلک

منفی	مثبت	
۳۲۱۰۹	۳۲۱۰۹	مدل با شکلک
۳۲۱۰۹	۳۲۱۰۹	مدل پایه

۴-۲- پیش‌پردازش داده

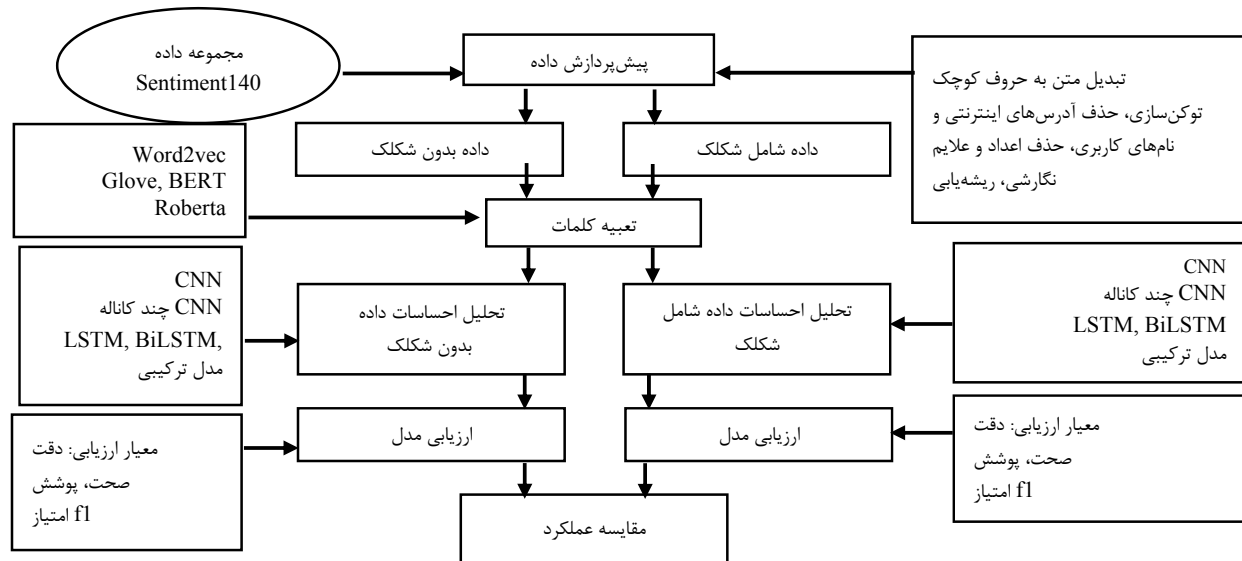
توییت‌های خام حاوی اطلاعات بی‌معنی و بدون ساختار و کلمات تکراری بودند. از این رو، هر آزمایش چند مرحله پیش‌پردازش را دنبال می‌کند. مجموعه داده‌ها مراحل پیش‌پردازش زیر را طی می‌کنند:

- ۱) تبدیل متن به حروف کوچک
 - ۲) حذف تمام فضاها و خالی اضافی
 - ۳) حذف همه اعداد و علائم نگارشی
 - ۴) حذف آدرس‌های اینترنتی، تگ‌های html و نام‌های کاربری
- برای مدیریت کلمات بن در ابتدا هر توییت را با استفاده از ابزار توکن‌ساز مجموعه ابزار زبان طبیعی^{۲۹} به لیستی از کلمات تقسیم می‌شود. برای ریشه یابی از Stemming استفاده می‌گردد. اگر هر یک از توکن‌ها به کوله کلمات NLTK تعلق داشت، کلمه مناسب از بن NLTK برای آن توکن جایگزین می‌شود. عمل ریشه‌یابی این امکان را می‌دهد که فرم‌های مختلف یک کلمه به یک فرم واحد تبدیل گردد. با این کار تعداد ویژگی‌ها کمتر می‌شود و کامپیوتر می‌تواند شکل‌های مختلف یک کلمه را یکی در نظر بگیرد. همچنین کلماتی که آپاستروف می‌گیرند (I'm, It's و غیره)، آپاستروف آن‌ها حذف شدند، چون توکن‌ساز^{۳۰} آن‌ها را به عنوان یک توکن جداگانه در نظر می‌گیرد.

۳-۴- معماری کلی مدل

با استفاده از معیارهای ارزیابی اندازه‌گیری می‌شود. مراحل کلی مدل پیشنهادی در شکل ۳ نشان داده شده است.

در مدل ارایه شده، در ابتدا مجموعه داده sentiment140 پیش‌پردازش می‌شود، سپس بردارهای تعبیه از قبل آموزش داده شده Word2vec، glove، bert و Roberta روی مجموعه داده اعمال می‌شود. در ادامه داده‌ها به الگوریتم‌های شبکه عصبی کانولوشنی، شبکه عصبی کانولوشنی چندکاناله، حافظه طولانی کوتاه‌مدت، حافظه طولانی کوتاه‌مدت دوجته و مدل ترکیبی تزریق می‌گردد و در نهایت دقت هر مدل



شکل ۳- معماری کلی مدل

۴-۴- معیار ارزیابی

که Acc_i ، Rec_i و Pre_i به ترتیب دقت، فراخوانی و صحت برای کلاس i هستند. N تعداد کلاس‌ها در طبقه‌بندی احساسات است. TP مثبت واقعی، FP مثبت کاذب، FN منفی کاذب و TN منفی واقعی است. TP_i تعداد نمونه‌هایی در کلاس i هستند که پیش‌بینی شده‌اند در کلاس i هستند. TN_i تعداد نمونه‌هایی که در کلاس i نیستند و پیش‌بینی هم شده‌اند که در کلاس i نیستند. FP_i تعداد نمونه‌هایی که در کلاس i نیستند و پیش‌بینی شده‌اند که در کلاس i هستند. FN_i تعداد نمونه‌هایی که در کلاس i هستند و پیش‌بینی شده‌اند که در کلاس i نیستند.

برای ارزیابی طبقه‌بندی احساسات، عملکرد مدل با استفاده از دقت، صحت، فراخوانی یا پوشش و امتیاز $f1$ اندازه‌گیری شد. دقت، معیار اندازه‌گیری تمام کلاس‌های مشخص شده در برابر کلاس واقعی است. صحت کسری از نمونه مربوطه در برابر نمونه پیش‌بینی شده است. پوشش کسری از نمونه‌های مطلوب است که روی کل مقادیر نمونه‌های مطلوب پیش‌بینی شده است. امتیاز $f1$ به عنوان میانگین وزنی پوشش و صحت تعریف می‌شود. معادلات زیر دقت، فراخوانی و صحت و امتیاز $F1$ را توصیف می‌کنند:

$$Acc_i = \frac{TP_i + TN_i}{TP_i + TN_i + FP_i + FN_i} \quad (1)$$

$$Rec_i = \frac{TP_i}{TP_i + FN_i} \quad (2)$$

$$Pre_i = \frac{TP_i}{TP_i + FP_i} \quad (3)$$

$$F1 = \frac{2 \times Pre \times Rec}{Pre + Rec} \quad (4)$$

آموزش می‌بیند و در طیف گسترده‌ای از برنامه‌های مدل زبانی استفاده شده است. در این تحقیق از BERT-base استفاده می‌گردد.

رویکرد BERT از قبل آموزش دیده بهینه شده قوی: مدل RoBERTa توسعه‌ای از BERT است. BERT و RoBERTa تحت خانواده ترانسفورمرها قرار می‌گیرند که برای رسیدگی به مشکل وابستگی‌های بلندمدت توسعه یافته است [۳۹]. RoBERTa همچنین روی مجموعه داده‌های بزرگ‌تری نسبت به BERT آموزش دیده است. این امر باعث می‌شود تا نمایش‌های RoBERTa در مقایسه با BERT بهتر تعمیم پیدا کنند. در این تحقیق از roberta-large-mnli استفاده می‌گردد.

در مرحله بعدی، لایه تعبیه از طریق چهار بردار تعبیه شده فوق آماده گردید. از تابع توکن‌ساز کراس استفاده شد که می‌تواند روی توییتهای مجموعه داده آموزشی برازش شود که متن را به رشته تبدیل می‌کند. این بردار نمایش کلمات به عنوان داده ورودی برای مدل‌های CNN و LSTM استفاده می‌شود که به نوبه خود برای طبقه‌بندی احساسات بکار برده می‌شود.

۴-۵- مدل شبکه عصبی کانولوشنی

با چند تغییر در مدل شبکه عصبی کانولوشنی کیم و با تنظیم پارامترها دو نوع CNN، ایجاد می‌گردد: یک مدل ساده CNN و یک مدل چندکاناله CNN. یک مدل ترتیبی CNN، دارای یک پشته خطی از لایه‌ها است که در شکل ۴ نشان داده شده است. حداکثر طول توییت ۴۶ بود، حداکثر طول در این دو مجموعه داده توسط توکن‌ساز کراس کنترل می‌شود. توکن‌ساز توییت‌ها را بردار می‌کند و آن‌ها را به دنباله‌ای از اعداد صحیح تبدیل می‌کند، سپس توکن‌ساز برای استفاده از رایج‌ترین کلمات محدود می‌شود. عمل پد گذاری با اضافه کردن مقدار ۰ در انتها، طول رشته‌های متنی را برابر با حداکثر طول یعنی طول بزرگ‌ترین رشته متنی می‌رساند. اولین لایه در این CNN یک لایه تعبیه^{۳۴} است که در قسمت لایه تعبیه قبل ساختیم. لایه تعبیه، خروجی را به یک لایه کانولوشنی^{۳۵} منتقل می‌کند. این لایه ماتریس خروجی را از طریق چند پارامتر مانند تعداد فیلترهای خروجی در کانولوشن‌ها، طول پنجره‌های کانولوشنی، تابع فعال‌سازی و تابع زیان از طریق روش میزان‌سازی^{۳۶} تولید می‌کند. از فیلترهایی استفاده می‌شود که در محدوده ۳-۸، اندازه هسته ۵-۳، تابع فعال‌سازی 'relu' یا 'tanh' و میزان‌سازی با وزن ۰/۰۰۰۲ هستند. خروجی لایه کانولوشنی به یک لایه حذف تصادفی با نرخ حذف ۵/۰ وارد می‌شود و به‌طور تصادفی به ۵۰ درصد نورون‌های شبکه وزن ۰ اختصاص می‌دهد. این عملیات حذف تصادفی شبکه باعث می‌شود حساسیت شبکه در هنگام واکنش به تغییرات کوچک‌تر در توییت کمتر شود که باعث افزایش دقت مدل می‌شود. سپس خروجی لایه حذف تصادفی به یک لایه کانولوشنی دوم وارد می‌شود.

لایه جاسازی یا تعبیه: قبل از ساختن مدل تجزیه و تحلیل احساسات، همه ویژگی‌های موجود در توییت‌ها باید استخراج شده و از طریق تعبیه کلمه در گروهی از کلمات مرتبط با معنا شکل بگیرند. این مطالعه از چند روش مختلف تعبیه کلمات، از جمله بردار تعبیه word2vec از قبل آموزش دیده، بردار تعبیه شده GloVe از قبل آموزش دیده، بردار تعبیه شده BERT از قبل آموزش داده شده و بردار تعبیه شده roberta از قبل آموزش داده شده استفاده می‌کند.

word2vec از قبل آموزش داده شده: برای طبقه‌بندی احساسات از بردار تعبیه word2vec از قبل آموزش دیده گوگل استفاده می‌گردد. دارای ۳۰۰ بعد بردار کلمات از قبل آموزش دیده است که با استفاده از ۱۰۰ میلیارد کلمه از یک مجموعه داده اخبار گوگل آموزش داده شده است. این بردار دارای اندازه واژگان حدود ۳ میلیون کلمه از اخبار گوگل است و به یک فرهنگ لغت با کلمات به عنوان کلید و ضرایب به عنوان مقادیر برای لایه تعبیه تبدیل شده است.

GloVe از قبل آموزش داده شده: برای طبقه‌بندی احساسات از بردار تعبیه GloVe از قبل آموزش دیده استنفورد استفاده می‌گردد. ۲۰۰ بعد از بردارهای کلمه از قبل آموزش دیده است که با استفاده از ۲ میلیارد توییت، ۲۷ میلیارد کلمه ساخته شده است و اندازه واژگان آن ۱،۲ میلیون است. این بردار به یک فرهنگ لغت با کلمات به عنوان کلید و ضرایب به عنوان مقادیر تبدیل شده است.

نمایش رمزگذار دوطرفه از ترانسفورمرها: نمایش رمزگذار دوطرفه از ترانسفورمرها توسط دولین^{۳۱} و همکاران (۲۰۱۸) مطرح شد که از رمزگذارها در یک ترانسفورمر به عنوان زیرساختی برای مدل‌های از قبل آموزش دیده برای کارهای پردازش زبان طبیعی مانند تحلیل احساسات، خلاصه‌سازی متن و غیره استفاده می‌کند [۳۸]. توانایی BERT برای وارد کردن دو جمله به‌طور هم‌زمان و تعیین اینکه آیا جمله دوم بعد از جمله اول می‌آید یا خیر، برای پیش‌بینی جمله بعد^{۳۲} (NSP) بسیار پرکاربرد می‌باشد. این توانایی به مدل کمک می‌کند تا روابط طولانی بین متون را حفظ کند. BERT دو مدل دارد، یعنی مدل‌های BERT-base و BERT-large. مدل BERT-base شامل بلوک‌های رمزگذار ترانسفورمر ۱۲ لایه است که هر بلوک حاوی ۱۲ لایه خود توجه^{۳۳} و ۷۶۸ لایه پنهان است و در مجموع ۱۱۰ میلیون پارامتر تولید می‌کند. از سوی دیگر، BERT-large از بلوک‌های رمزگذار ترانسفورمر ۲۴ لایه تشکیل شده است که هر بلوک حاوی ۲۴ لایه خود توجه است و در مجموع ۳۴۰ میلیون پارامتر تولید می‌کند. عملکرد BERT به نوع مدل بستگی دارد، BERT-large می‌تواند دقت بالاتری نسبت به BERT-base داشته باشد. با این حال، این بهبود در دقت با استفاده از BERT-large به منابع بیشتری نیاز دارد. مزیت BERT شامل توانایی آن در مدیریت استخراج اطلاعات متنی به دلیل توانایی دو جهته آن است. سریع‌تر

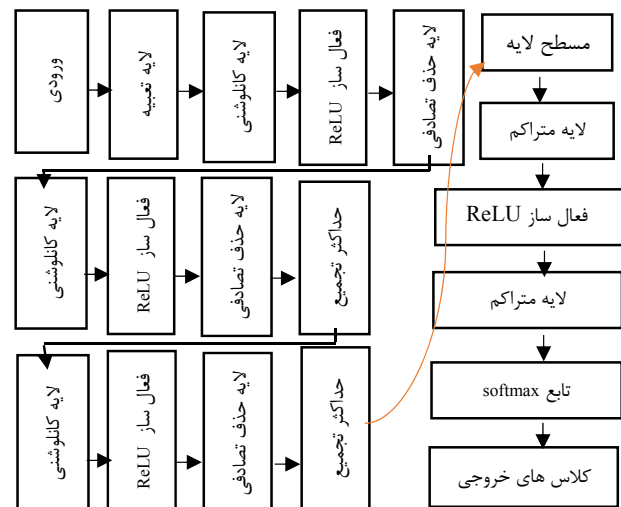
این مدل دارای سه کانال است. هر کانال دارای یک لایه ورودی است. لایه ورودی به یک لایه تعبیه، یک لایه یک‌بعدی کانولوشنی، یک لایه حذف تصادفی، یک لایه حداکثر تجمیع و یک لایه Flatten ارسال می‌شود. بردارهای خروجی لایه Flatten هر سه کانال به هم متصل شده و به یک لایه متراکم با تابع فعال‌سازی 'relu' یا 'tanh' تغذیه می‌شود. خروجی لایه فعال‌سازی به یک لایه متراکم با تابع فعال‌سازی softmax منتقل می‌شود. این دو مدل CNN با استفاده از تابع کراس categorical_crossentropy به‌عنوان تابع هزینه و "adam" به‌عنوان تابع بهینه‌ساز از طریق تابع کامپایل کراس کامپایل شدند. انواع مدل‌های ساده و چندکاناله CNN در جدول ۴ نشان داده شده است.

جدول ۴- مدل‌های مختلف شبکه عصبی کانولوشنی

مدل	Word2vec	GloVe	BERT
Cnn ساده	✓	✓	✓
Cnn چندکاناله	✓	✓	✓

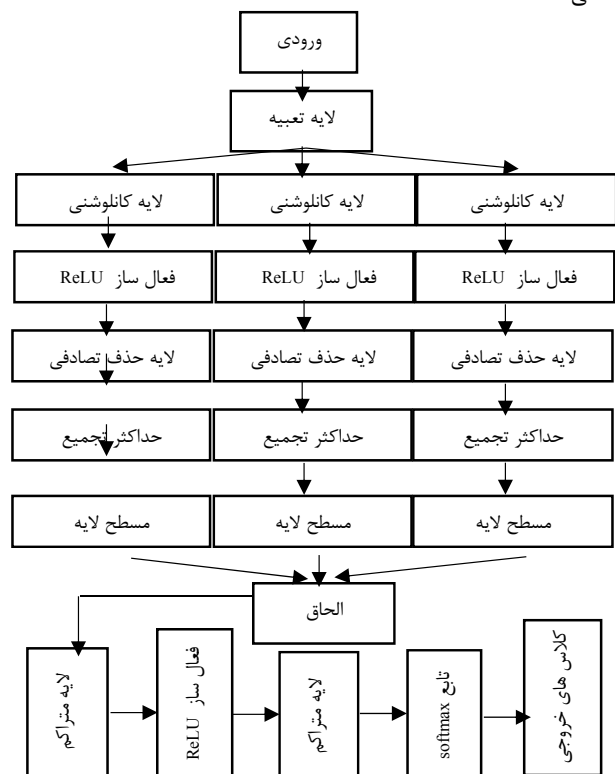
۴-۶- مدل حافظه طولانی کوتاه‌مدت

این تحقیق از یک معماری برای هر دو مدل حافظه طولانی کوتاه‌مدت و حافظه طولانی کوتاه‌مدت دو جهته استفاده می‌کند که از معماری مدل BiLSTM کلیشی^{۲۸} پیروی می‌کند که در شکل ۶ نشان داده شده است. این مدل دو LSTM برای آموزش در توییت‌های ورودی ایجاد می‌کند. طول توییت‌های مجموعه داده آموزشی و تست برای این مدل LSTM در همه توییت‌ها نرمالایز شد. نرمال‌سازی متن ورودی مشابه مدل CNN بود که در قسمت قبل توضیح داده شد. شبکه LSTM یک پشته خطی از لایه‌ها است که با استفاده از کراس ساخته شده است. لایه اول یک لایه تعبیه است که رشته ورودی را به رشته‌ای از بردارهای عددی کد می‌کند. خروجی لایه تعبیه به یک لایه LSTM متصل است. LSTM رشته بردارها را بر اساس پارامتر اندازه بعد خروجی که حاوی اطلاعات مربوط به کل رشته است، به یک بردار واحد تبدیل می‌کند. لایه LSTM دوطرفه از طریق کراس پیاده‌سازی شد. خروجی لایه LSTM به یک لایه متراکم منتقل شده است و تابع فعال‌سازی این لایه "relu" یا "tanh" است. خروجی این لایه به یک لایه متراکم می‌رود که از softmax به‌عنوان تابع فعال‌سازی استفاده می‌کند، که باعث ایجاد برچسب احساسات شد.



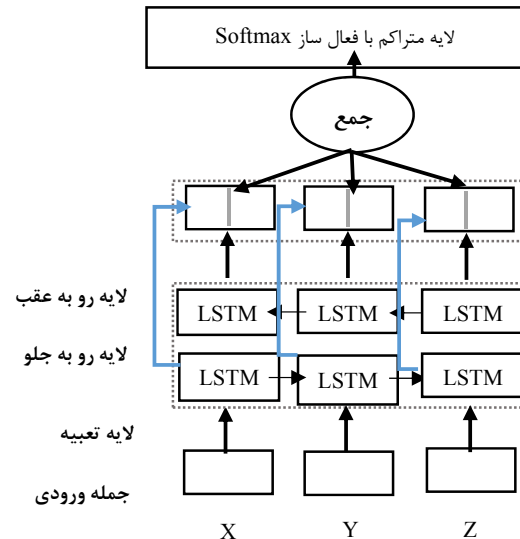
شکل ۴- مدل شبکه عصبی کانولوشنی ساده

سیس این لایه کانولوشنی به یک لایه حذف تصادفی با نرخ حذف ۰/۳ منتقل شد. به همین طریق، لایه حذف تصادفی به لایه حداکثر تجمیع و لایه حداکثر تجمیع به لایه کانولوشنی سوم منتقل شد. این لایه کانولوشنی سوم به یک لایه حذف تصادفی با نرخ حذف تصادفی ۰/۱ و این لایه حذف تصادفی به یک لایه حداکثر تجمیع منتقل شد. شکل ۵ مدل CNN چندکاناله را که در این مقاله استفاده شده است به تصویر می‌کشد.



شکل ۵- شبکه عصبی کانولوشنی چندکاناله

توجه بر روی خروجی لایه‌های BiLSTM برای تأکید بیشتر بر کلمات مهم‌تر و تأکید کمتر بر کلمات کم‌اهمیت‌تر اعمال می‌شود [۴۰]. خروجی لایه توجه به یک لایه متراکم منتقل می‌شود و تابع فعال‌سازی این لایه "relu" است. در نهایت یک لایه متراکم مورد استفاده قرار می‌گیرد که از softmax به عنوان تابع فعال‌سازی استفاده می‌کند؛ که باعث ایجاد برچسب احساسات می‌شود. مدل پیشنهادی در شکل ۷ نشان داده شده است. در توییت بعضی از کلمات ممکن است عبارات خاص را نادرست طبقه‌بندی کند، برای این منظور یک لایه حذف تصادفی جهت جلوگیری از پیش‌برازش به شبکه اضافه شد، که خروجی لایه کانولوشن را دریافت می‌کند و کلمات را تصادفی خارج می‌کند. همچنین یک مدل CNN چند کاناله شامل استفاده از چندین ورژن از مدل CNN با اندازه هسته‌های مختلف ایجاد می‌گردد. لایه حذف تصادفی محبوب‌ترین رویکرد برای میزان‌سازی CNN و کاهش مسائل پیش‌برازش است. این لایه از انطباق هم‌زمان سلول‌های عصبی جلوگیری کرده و آن‌ها را مجبور به یادگیری ویژگی‌های ارزشمند می‌کند. این مدل CNN از "adagrad" یا "adam" به عنوان بهینه‌ساز استفاده می‌کند.



شکل ۶- مدل BiLSTM

هر دو آزمایش با مدل LSTM و مدل BiLSTM، همراه با چهار مدل مختلف تعبیه که در جدول ۵ نشان داده شده است، مورد بررسی قرار گرفت.

جدول ۵- مدل‌های حافظه طولانی کوتاه مدت

مدل	Word2vec	GloVe	BERT	roberta
LSTM	✓	✓	✓	✓
BiLSTM	✓	✓	✓	✓

۴-۸- نتایج مدل‌های CNN

مدل‌های ساده و چندکاناله CNN برای طبقه‌بندی احساسات مورد آزمایش قرار گرفتند. این مدل‌ها با استفاده از سه بردار تعبیه مختلف برای مجموعه داده‌های مدل پایه و مدل شامل شکلک آموزش و ارزیابی شدند. در مدل‌های CNN بر اساس زمان اجرا و عملکرد مدل در هر مجموعه داده چند پارامتر تنظیم گردید.

۴-۸-۱- CNN با word2vec، Glove و BERT از قبل آموزش دیده

سه مدل CNN پیاده‌سازی شد که در هر کدام یک لایه تعبیه با استفاده از یک مورد از بردارهای word2vec، Glove و BERT از قبل آموزش داده شده ایجاد گردید. این مدل‌ها با استفاده از هر دو مجموعه داده آموزشی آموزش داده شدند. سپس دقت پیش‌بینی مجموعه داده تست برای هر دو مجموعه داده اندازه‌گیری شد. جدول ۶ دقت هر کدام از مدل‌ها را نشان می‌دهد.

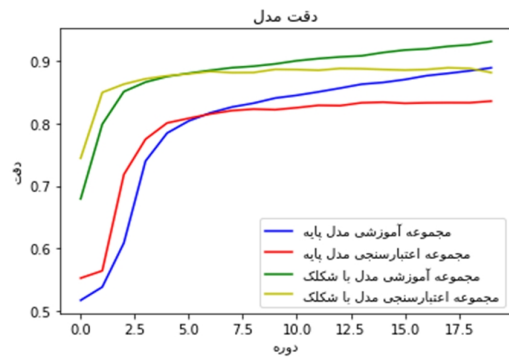
جدول ۶- دقت مدل‌های CNN روی مجموعه داده پایه و شامل شکلک

مدل	دقت (Accuracy) (%)	
	مدل پایه	مدل با شکلک
word2vec + CNN	۷۸/۹۶	۸۴/۸۷
GloVe + CNN	۸۰/۱۶	۸۵/۶۶
BERT + CNN	۸۲/۹۲	۸۷/۶۶

نتایج نشان می‌دهد مدل CNN با بردار تعبیه BERT از قبل آموزش داده شده نسبت به دو مدل دیگر عملکرد بهتری دارد.

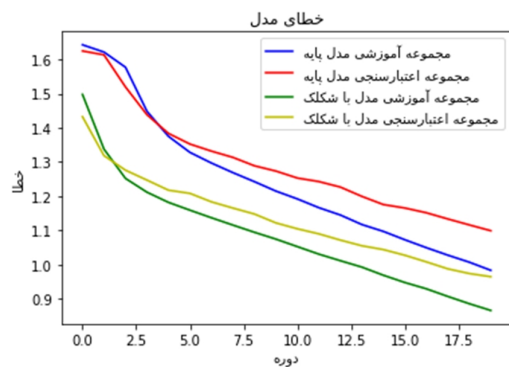
۴-۷- مدل ترکیبی مبتنی بر توجه شبکه‌های عصبی کانولوشنی و شبکه‌های حافظه طولانی کوتاه مدت دو جهته (ACBR۴۸)

معماری پیشنهادی از هفت بخش اصلی تشکیل شده است، که شامل پیش‌پردازش داده، جاسازی کلمه، لایه کانولوشنی، لایه تجمیع، لایه‌های شبکه حافظه طولانی کوتاه مدت دوطرفه، و دو لایه دیگر می‌باشد (شکل ۷). در لایه جاسازی برای بدست آوردن نمایش برداری نظرات متنی و ایجاد بردار ویژگی‌ها RoBERTa مورد استفاده قرار می‌گیرد. هدف RoBERTa ایجاد یک جاسازی کلمه معنی‌دار به عنوان نمایش ویژگی است، به طوری که لایه‌های بعدی به راحتی می‌توانند اطلاعات مفید را از جاسازی کلمه دریافت کنند. خروجی لایه جاسازی به عنوان ورودی به لایه کانولوشنی وارد می‌شود که یک نگاشت ویژگی ایجاد می‌کند که حضور ویژگی‌های شناخته شده در ورودی را خلاصه می‌کند و ویژگی‌های ذاتی و محلی را بدست می‌آورد. سپس لایه تجمیع جهت کاهش ابعاد نگاشت ویژگی مورد استفاده قرار می‌گیرد. همچنین یک مکانیسم



شکل ۸ - CNN با BERT از قبل آموزش دیده - دقت مدل

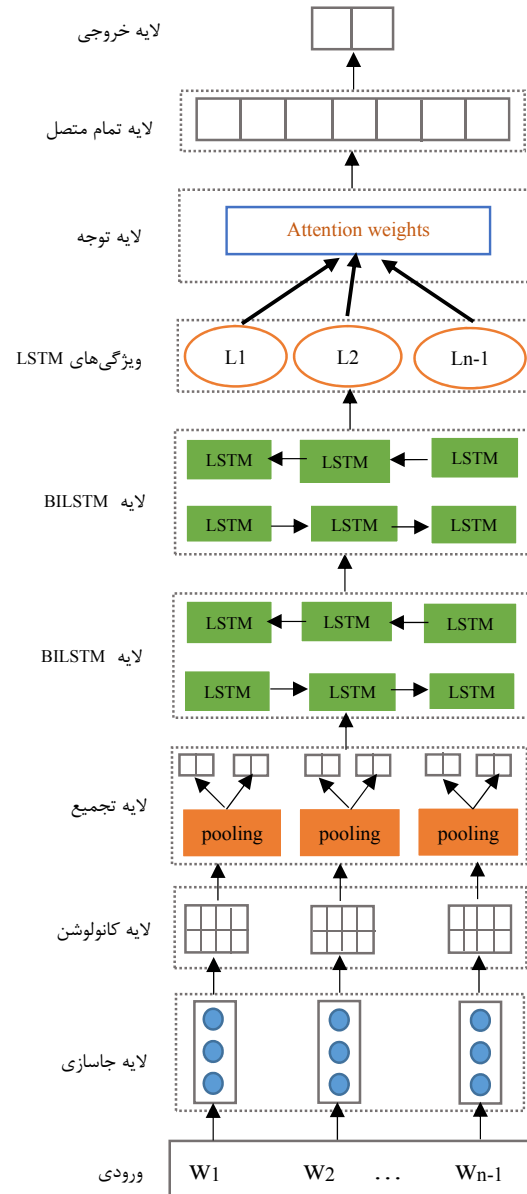
شکل ۸، همگرایی مدل را نشان می‌دهد که بر اساس آزمایش‌های انجام شده بر روی مجموعه داده‌های آموزشی با تکرار ۲۰ دور بدست آمده است. دقت بدست آمده ۸۷/۸۰ درصد می‌باشد که برای تکرارهای بعدی نیز تقریباً ثابت می‌ماند.



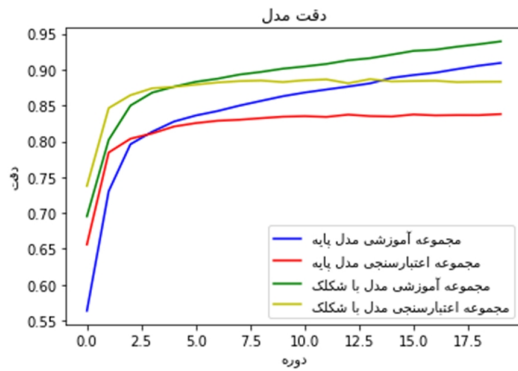
شکل ۹ - CNN با BERT از قبل آموزش دیده - خطای مدل

جدول ۷ - CNN با BERT از قبل آموزش دیده - صحت، پوشش، امتیاز f1 (مدل پایه)

	مدل پایه			
	صحت	پوشش	معیار f1	پشتیبانی
۰	۸۵	۸۲	۸۳	۶۶۷۴
۱	۸۱	۸۴	۸۲	۶۱۷۰
دقت			۸۳	۱۲۸۴۴
میانگین کلان	۸۳	۸۳	۸۳	۱۲۸۴۴
میانگین وزنی	۸۳	۸۳	۸۳	۱۲۸۴۴

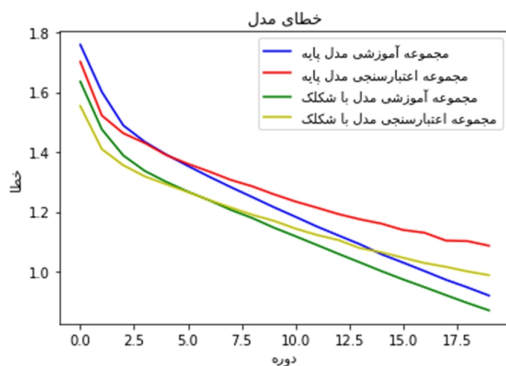


شکل ۷- مدل ترکیبی مبتنی بر توجه شبکه‌های عصبی کانولوشنی و شبکه‌های حافظه طولانی کوتاه مدت دوجهته
نتایج نشان می‌دهد مدل CNN با بردار تعبیه BERT از قبل آموزش داده شده نسبت به دو مدل دیگر عملکرد بهتری دارد. شکل ۸ و شکل ۹ دقت و خطا را برای این مدل نشان می‌دهد. دقت برای مدل پایه و مدل با شکلک به ترتیب ۸۲/۹۲ و ۸۷/۶۶ درصد است. جدول ۷ و جدول ۸ صحت، پوشش و امتیاز f1 را برای مجموعه داده تست مدل پایه و مدل شامل شکلک شرح می‌دهد.



شکل ۱۰ - CNN چندکاناله با BERT از قبل آموزش دیده - دقت مدل

شکل ۱۰، همگرایی مدل را نشان می‌دهد که بر اساس آزمایش‌های انجام شده بر روی مجموعه داده‌های آموزشی با تکرار ۲۰ دور بدست آمده است. دقت بدست آمده ۸۸/۲۰ درصد می‌باشد که برای تکرارهای بعدی نیز تقریباً ثابت می‌ماند.



شکل ۱۱ - CNN چندکاناله با BERT از قبل آموزش دیده - خطای مدل

جدول ۱۰ - CNN چندکاناله با BERT از قبل آموزش دیده - صحت، پوشش، امتیاز f1 (مدل پایه)

	مدل پایه			
	پشتیبانی	معیار f1	پوشش	صحت
۰	۶۶۱۶	۸۳	۸۳	۸۴
۱	۶۲۲۸	۸۳	۸۴	۸۲
دقت	۱۲۸۴۴	۸۳		
میانگین کلان	۱۲۸۴۴	۸۳	۸۳	۸۳
میانگین وزنی	۱۲۸۴۴	۸۳	۸۳	۸۳

جدول ۱۱ - CNN چندکاناله با BERT از قبل آموزش دیده - صحت، پوشش، امتیاز f1 (مدل شامل شکلک)

مدل با شکلک

جدول ۸ - CNN با BERT از قبل آموزش دیده - صحت، پوشش، امتیاز

f1 (مدل شامل شکلک)

مدل با شکلک				
پشتیبانی	معیار f1	پوشش	صحت	
۶۶۲۲	۸۸	۸۷	۸۹	۰
۶۲۲۱	۸۷	۸۸	۸۶	۱
۱۲۸۴۴	۸۸			دقت
۱۲۸۴۴	۸۸	۸۸	۸۸	میانگین کلان
۱۲۸۴۴	۸۸	۸۸	۸۸	میانگین وزنی

۴-۸-۲- چندکاناله با word2vec, Glove و BERT از قبل آموزش دیده

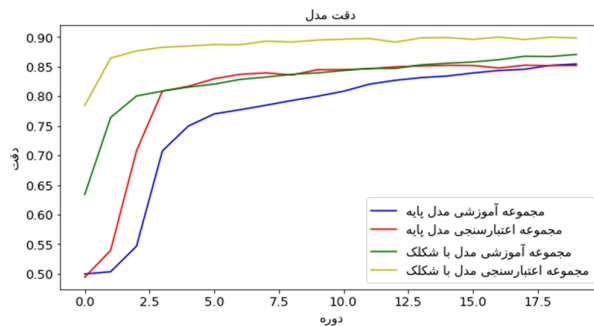
سه مدل CNN چندکاناله پیاده‌سازی شد که در هر کدام یک لایه تعبیه با استفاده از یک مورد از بردارهای Glove, word2vec و BERT از قبل آموزش داده شده ایجاد گردید. این مدل‌ها با استفاده از هر دو مجموعه داده آموزشی آموزش داده شدند سپس دقت پیش‌بینی مجموعه داده تست برای هر دو مجموعه داده اندازه‌گیری شد. جدول ۹ دقت هر کدام از مدل‌ها را نشان می‌دهد.

جدول ۹ - دقت مدل‌های CNN چندکاناله روی مجموعه داده پایه و

مجموعه داده شامل شکلک

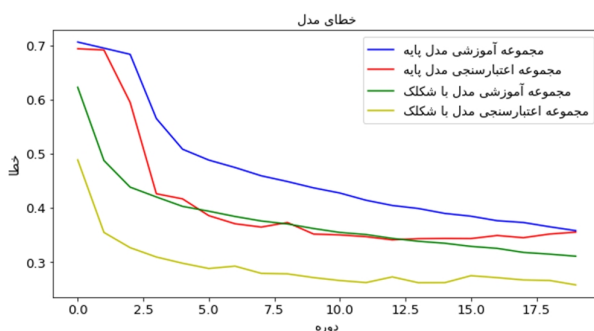
مدل	دقت (Accuracy) (%)	
	مدل پایه	مدل با شکلک
CNN چندکاناله + word2vec	۷۹/۷۲	۸۵/۰۳
CNN چندکاناله + GloVe	۸۱/۰۸	۸۷/۴۲
CNN چندکاناله + BERT	۸۳/۰۶	۸۸/۰۳

نتایج نشان می‌دهد مدل CNN چندکاناله با بردار تعبیه BERT از قبل آموزش داده شده نسبت به دو مدل دیگر عملکرد بهتری دارد. شکل ۱۰ و شکل ۱۱ دقت و خطا را برای این مدل نشان می‌دهد. دقت برای مدل پایه و مدل شامل شکلک به ترتیب ۸۳/۰۶ و ۸۸/۰۳ درصد است. جدول ۱۰ و جدول ۱۱ صحت، پوشش و امتیاز f1 را برای مجموعه داده تست مدل پایه و مدل شامل شکلک نشان می‌دهند.



شکل ۱۲ - LSTM با RoBERTa از قبل آموزش دیده - دقت مدل

شکل ۱۲، همگرایی مدل را نشان می‌دهد که بر اساس آزمایش‌های انجام شده بر روی مجموعه داده‌های آموزشی با تکرار ۲۰ دور بدست آمده است. دقت بدست آمده ۸۹/۴۸ درصد می‌باشد که برای تکرارهای بعدی نیز تقریباً ثابت می‌ماند. شکل ۱۳ خطای مدل را نشان می‌دهد که که در تکرار ۱۴ به ۲۶/۲۳ درصد کاهش می‌یابد.



شکل ۱۳ - LSTM با RoBERTa از قبل آموزش دیده - خطای مدل

جدول ۱۳ - LSTM با RoBERTa از قبل آموزش دیده - صحت، پوشش، امتیاز f1 (مدل پایه)

مدل پایه				
	پشتیبانی	معیار f1	پوشش	صحت
۰	۶۶۸۲	۸۶	۸۵	۸۷
۱	۶۱۸۲	۸۵	۸۶	۸۴
دقت	۱۲۸۴۴	۸۵		
میانگین کلان	۱۲۸۴۴	۸۵	۸۵	۸۵
میانگین وزنی	۱۲۸۴۴	۸۵	۸۵	۸۵

پشتیبانی	معیار f1	پوشش	صحت	
۶۷۹۰	۸۸	۸۶	۹۰	۰
۶۰۵۴	۸۸	۹۰	۸۶	۱
۱۲۸۴۴	۸۸			دقت
۱۲۸۴۴	۸۸	۸۸	۸۸	میانگین کلان
۱۲۸۴۴	۸۸	۸۸	۸۸	میانگین وزنی

۵- نتایج مدل‌های LSTM

یک مدل LSTM و یک مدل LSTM دوطرفه (BiLSTM) برای طبقه‌بندی احساسات ایجاد گردید. هر دو مدل آموزش داده شده و با استفاده از چهار بردار تعبیه شده برای مجموعه داده پایه و مجموعه داده شامل شکلک پیش‌بینی شده‌اند.

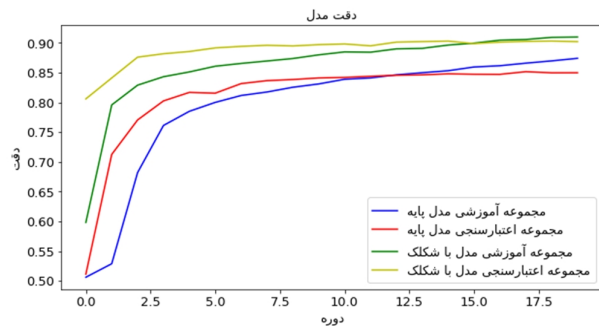
۵-۱- LSTM با word2vec، GloVe، BERT و RoBERTa از قبل آموزش دیده

چهار مدل LSTM پیاده‌سازی شد که در هر کدام یک لایه تعبیه با استفاده از یک مورد از بردارهای word2vec، GloVe، BERT و RoBERTa از قبل آموزش داده شده ایجاد گردید. این مدل‌ها با استفاده از هر دو مجموعه داده آموزشی آموزش داده شدند سپس دقت پیش‌بینی مجموعه داده تست برای هر دو مجموعه داده اندازه‌گیری شد. جدول ۱۲ دقت هر کدام از مدل‌ها را نشان می‌دهد.

جدول ۱۲- دقت مدل‌های LSTM روی مجموعه داده پایه و مجموعه داده شامل شکلک

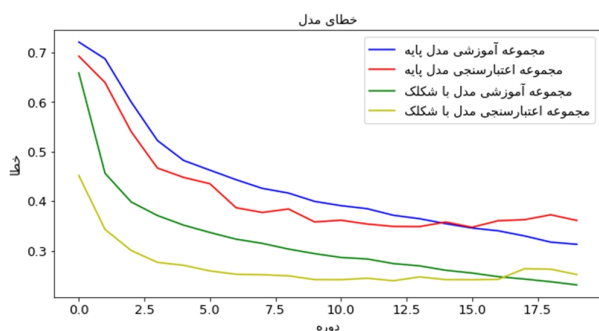
مدل	دقت (Accuracy) (%)	
	مدل پایه	مدل با شکلک
word2vec + LSTM	۸۰/۴۹	۸۶/۹۰
GloVe + LSTM	۸۰/۹۱	۸۷/۱۶
BERT + LSTM	۸۲/۷۰	۸۸/۵۰
RoBERTa + LSTM	۸۵/۳۰	۸۹/۳۷

نتایج نشان می‌دهد مدل LSTM با بردار تعبیه RoBERTa از قبل آموزش داده شده نسبت به سه مدل دیگر عملکرد بهتری دارد. شکل ۱۲ و شکل ۱۳ دقت و خطا را برای این مدل نشان می‌دهد. دقت برای مدل پایه و مدل شامل شکلک به ترتیب ۸۵/۳۰ و ۸۹/۳۷ درصد است. جدول ۱۳ و جدول ۱۴ صحت، پوشش و امتیاز f1 را برای مجموعه داده تست مدل پایه و مدل شامل شکلک شرح می‌دهد.



شکل ۱۴- BiLSTM با RoBERTa از قبل آموزش دیده - دقت مدل

شکل ۱۴، همگرایی مدل را نشان می‌دهد که بر اساس آزمایش‌های انجام شده بر روی مجموعه داده‌های آموزشی با تکرار ۲۰ دور بدست آمده است. دقت بدست آمده ۸۹/۵۲ درصد می‌باشد که برای تکرارهای بعدی نیز تقریباً ثابت می‌ماند. شکل ۱۵ خطای مدل را نشان می‌دهد که که در تکرار ۱۵ به ۲۳/۹۶ درصد کاهش می‌یابد.



شکل ۱۵- BiLSTM با RoBERTa از قبل آموزش دیده - خطای مدل

جدول ۱۵ - BiLSTM با RoBERTa از قبل آموزش دیده - صحت، پوشش، امتیاز f1 (مدل پایه)

مدل پایه				
	صحت	پوشش	معیار f1	پشتیبانی
۰	۸۶	۸۵	۸۶	۶۴۸۹
۱	۸۵	۸۵	۸۵	۶۳۵۵
دقت			۸۵	۱۲۸۴۴
میانگین کلان	۸۵	۸۵	۸۵	۱۲۸۴۴
میانگین وزنی	۸۵	۸۵	۸۵	۱۲۸۴۴

جدول ۱۴ - LSTM با RoBERTa از قبل آموزش دیده - صحت، پوشش، امتیاز f1 (مدل شامل شکلک)

مدل با شکلک				
	صحت	پوشش	معیار f1	پشتیبانی
۰	۹۲	۸۸	۹۰	۶۷۴۶
۱	۸۷	۹۱	۸۹	۶۰۸۰
دقت			۸۹	۱۲۸۴۴
میانگین کلان	۸۹	۸۹	۸۹	۱۲۸۴۴
میانگین وزنی	۸۹	۸۹	۸۹	۱۲۸۴۴

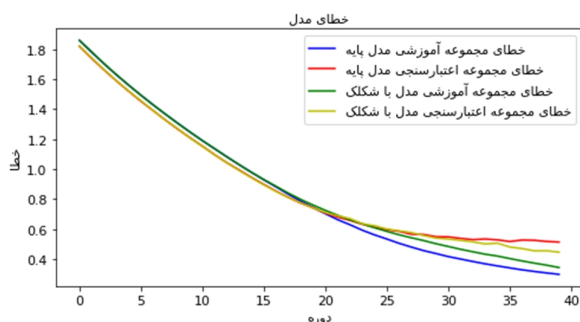
۵-۲- BiLSTM با word2vec، Glove، BERT و RoBERTa از قبل آموزش دیده

چهار مدل BiLSTM پیاده‌سازی شد که در هر کدام یک لایه تعبیه با استفاده از یک مورد از بردارهای word2vec، Glove، BERT و RoBERTa از قبل آموزش داده شده ایجاد گردید. این مدل‌ها با استفاده از هر دو مجموعه داده آموزشی آموزش داده شدند، سپس دقت پیش-بینی مجموعه داده تست برای هر دو مجموعه داده اندازه‌گیری شد. جدول ۱۵ دقت هر کدام از مدل‌ها را نشان می‌دهد.

جدول ۱۵- دقت مدل‌های BiLSTM روی مجموعه داده پایه و مجموعه داده شامل شکلک

مدل	دقت (Accuracy) (%)	
	مدل پایه	مدل با شکلک
word2vec + BiLSTM	۸۱/۲۱	۸۷/۲۴
GloVe + BiLSTM	۸۱/۲۸	۸۷/۰۱
BERT + BiLSTM	۸۳/۳۲	۸۸/۵۳
RoBERTa + BiLSTM	۸۵/۳۷	۸۹/۵۹

نتایج نشان می‌دهد مدل BiLSTM با بردار تعبیه RoBERTa از قبل آموزش داده شده نسبت به سه مدل دیگر عملکرد بهتری دارد. شکل ۱۴ و شکل ۱۵ دقت و خطا را برای این مدل نشان می‌دهد. دقت برای مدل پایه و مدل با شکلک به ترتیب ۸۵/۳۰ و ۸۹/۳۷ درصد است. جدول ۱۵ و جدول ۱۶ صحت، پوشش و امتیاز f1 را برای مجموعه داده تست مدل پایه و مدل شامل شکلک شرح می‌دهد.



شکل ۱۷- ACBR با RoBERTa از قبل آموزش دیده - خطای مدل

در ارزیابی نتایج شکل‌های ۱۶ و ۱۷، دقت و خطای مدل ACBR پیشنهادی پس از تعداد دوره‌های ذکر شده همگرا می‌شود، بنابراین، این تحلیل‌ها بیانگر این است که مدل ACBR مشکل پیش برآزش را کاهش داده و دقت را افزایش می‌دهد.

جدول ۱۷ - ACBR با RoBERTa از قبل آموزش دیده - صحت، پوشش، امتیاز f1 (مدل پایه)

	مدل پایه			پشتیبانی
	صحت	پوشش	معیار f1	
۰	۸۶	۸۶	۸۶	۶۵۱۳
۱	۸۵	۸۶	۸۶	۶۳۳۱
دقت			۸۶	۱۲۸۴۴
میانگین کلان	۸۶	۸۶	۸۶	۱۲۸۴۴
میانگین وزنی	۸۶	۸۶	۸۶	۱۲۸۴۴

جدول ۱۸ - ACBR با RoBERTa از قبل آموزش دیده - صحت، پوشش، امتیاز f1 (مدل شامل شکلک)

	مدل با شکلک			پشتیبان
	صحت	پوشش	معیار f1	
۰	۹۲	۸۹	۹۰	۶۷۱۹
۱	۸۸	۹۲	۹۰	۶۱۲۵
دقت			۹۰	۱۲۸۴۴
میانگین کلان	۹۰	۹۰	۹۰	۱۲۸۴۴
میانگین وزنی	۹۰	۹۰	۹۰	۱۲۸۴۴

۶- ارزیابی مدل‌ها

همه مدل‌های CNN و LSTM با استفاده از مجموعه داده‌های آموزشی آزمایش‌های مدل پایه و مدل شامل شکلک با پارامترهایی نشان داده شده در جداول ۱۹ و ۲۰ آموزش داده شدند. در حین برآزش هر مدل، عملکرد

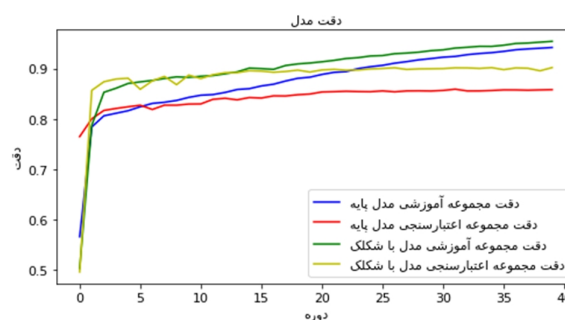
جدول ۱۶ - BiLSTM با RoBERTa از قبل آموزش دیده - صحت، پوشش، امتیاز f1 (مدل شامل شکلک)

مدل با شکلک				
	صحت	پوشش	معیار f1	پشتیبانی
۰	۹۲	۸۸	۹۰	۶۷۳۹
۱	۸۷	۹۱	۸۹	۶۱۰۵
دقت			۹۰	۱۲۸۴۴
میانگین کلان	۹۰	۹۰	۹۰	۱۲۸۴۴
میانگین وزنی	۹۰	۹۰	۹۰	۱۲۸۴۴

۵-۳- نتایج مدل ACBR با roberta از قبل آموزش دیده

در این مدل، یک لایه تعبیه با استفاده از بردار RoBERTa از قبل آموزش داده شده برای مدل ACBR ایجاد گردید. این مدل با استفاده از هر دو مجموعه داده آموزشی آموزش داده شد و با استفاده از مجموعه داده اعتبارسنجی، دقت، صحت، پوشش، امتیاز f1 و خطای مدل برای هر دو آزمایش اندازه‌گیری شد.

شکل ۱۶ و شکل ۱۷ این معیارها را برای این مدل نشان می‌دهد. سپس دقت پیش‌بینی مجموعه داده تست برای هر دو مجموعه داده اندازه‌گیری شد. دقت برای مدل پایه و مدل شامل شکلک به ترتیب ۸۵/۸۰ و ۹۰/۱۸ درصد است. جدول ۱۷ و جدول ۱۸ صحت، پوشش و امتیاز f1 را برای مجموعه داده تست مدل پایه و مدل شامل شکلک شرح می‌دهد.



شکل ۱۶- ACBR با RoBERTa از قبل آموزش دیده - دقت مدل

روی مجموعه داده پایه ۸۵/۸۰٪ و روی مجموعه داده شامل شکلک ۹۰/۱۸٪ است.

جدول ۲۱ - دقت مدل‌ها روی مجموعه داده پایه و مجموعه داده شامل شکلک

اختلاف (%)	دقت (Accuracy)		مدل
	مدل شامل شکلک	مدل پایه	
۵/۹۱	۸۴/۸۷	۷۸/۹۶	word2vec + CNN
۵/۵	۸۵/۶۶	۸۰/۱۶	GloVe + CNN
۴/۷۴	۸۷/۶۶	۸۲/۹۲	BERT + CNN
۵/۳۱	۸۵/۰۳	۷۹/۷۲	CNN چندکاناله + word2vec
۶/۳۴	۸۷/۴۲	۸۱/۰۸	CNN چندکاناله + GloVe
۴/۹۷	۸۸/۰۳	۸۳/۰۶	CNN چندکاناله + BERT
۶/۴۱	۸۶/۹۰	۸۰/۴۹	word2vec + LSTM
۶/۲۵	۸۷/۱۶	۸۰/۹۱	GloVe + LSTM
۵/۸	۸۸/۵۰	۸۲/۷۰	BERT + LSTM
۴/۰۷	۸۹/۳۷	۸۵/۳۰	roberta + LSTM
۶/۰۳	۸۷/۲۴	۸۱/۲۱	word2vec + BiLSTM
۵/۷۳	۸۷/۰۱	۸۱/۲۸	GloVe + BiLSTM
۵/۲۱	۸۸/۵۳	۸۳/۳۲	BERT + BiLSTM
۴/۲۲	۸۹/۵۹	۸۵/۳۷	roberta + BiLSTM
۴/۳۲	۹۰/۱۸	۸۵/۸۰	roberta + ACBR

نتایج نشان می‌دهد عملکرد همه مدل‌ها برای مجموعه داده شامل شکلک بهتر از مجموعه داده پایه است. مدل‌های LSTM و BiLSTM نسبت به مدل‌های CNN و CNN چندکاناله برای این طبقه‌بندی احساسات بهتر عمل کردند. BiLSTM نسبت به LSTM کمی بهتر عمل کرده است و مدل ACBR دارای بالاترین دقت در هر دو مجموعه داده است. همچنین بردارهای از قبل آموزش‌دیده roberta و BERT از بردارهای از قبل آموزش‌دیده GloVe و word2vec و بردار roberta از بردار BERT دقت بهتری را ارائه کردند.

برای دریافت بینش بهتر از بهترین رویکردهای کمی، صحت، پوشش، امتیاز f1، پشتیبانی، دقت و میانگین کلان در همه مدل‌ها در نظر گرفته شد. مدل ACBR با بردار تعبیه roberta از پیش آموزش‌دیده بهترین نتایج را در بین مدل‌های آزمایش‌شده به دست آورد. صحت، پوشش و امتیاز f1 برای مجموعه داده پایه ۸۶/۰، ۸۶/۰، ۸۶/۰ و برای مجموعه داده شامل شکلک به ترتیب ۹۰/۰، ۹۰/۰ و ۹۰/۰ بود. از این‌رو، مدل ACBR با بردار تعبیه roberta از پیش آموزش‌دیده به‌عنوان بهترین مدل برای مجموعه داده پایه و مجموعه داده شامل شکلک توصیه می‌شود.

مدل با استفاده از مجموعه داده اعتبارسنجی موردسنجش قرار گرفت. برای جلوگیری از پیش برآزش یک تنظیم‌کننده L2 با نرخ ۰/۰۰۰۲ مورد استفاده قرار گرفت، همچنین نرخ یادگیری کاهش یافت و لایه‌های حذف تصادفی به مدل‌ها اضافه گردید. عملکرد مدل با استفاده از دقت، صحت، پوشش و امتیاز F1 اندازه‌گیری شد. برای مدل‌های CNN و CNN چندکاناله از پارامترهای نشان داده‌شده در جدول ۱۹ استفاده گردید.

جدول ۱۹ - پارامترهای پیشنهادشده مدل‌های شبکه عصبی کانولوشنی

پارامتر	مقدار
حداکثر تعداد کلمات	۴۵-۴۶
تعداد کلاس‌های لایه خروجی	۲
اندازه واژه‌نامه	۴۹۲۲۷-۳۸۱۹۶
اندازه دسته	۱۲۸
توابع فعال‌سازی	relu
تعداد تکرار (چرخش)	۲۰
نرخ حذف تصادفی	۰/۱-۰/۳-۰/۵
بهینه‌ساز	Adam
تابع زیان	categorical_crossentropy
نوع Pooling	2,max

برای مدل‌های LSTM، BiLSTM و ACBR پارامترهای نشان داده‌شده در جدول ۲۰ مورد استفاده قرار گرفت.

جدول ۲۰ - پارامترهای پیشنهادشده مدل‌های حافظه طولانی کوتاه مدت

پارامتر	مقدار (LSTM)	مقدار (ACBR)
حداکثر تعداد کلمات	۴۵-۴۶	۴۵-۴۶
تعداد کلاس‌های لایه خروجی	۲	۲
اندازه واژه‌نامه	۴۹۲۲۷-۳۸۱۹۶	۴۹۲۲۷-۳۸۱۹۶
اندازه دسته	۱۲۸	۱۲۸
توابع فعال‌سازی	Relu	Relu
تعداد تکرار (چرخش)	۲۰	۴۰
نرخ حذف تصادفی	۰/۵	۰/۵
بهینه‌ساز	Adam	Adam
تابع زیان	categorical_crossentropy	Categorical_crossentropy
تعداد نودهای لایه متراکم	۱۰	۲۰

۶-۱- مقایسه

همه مدل‌ها با معیارهای عملکرد مختلف مانند دقت، صحت، پوشش و امتیاز f1 مقایسه شدند. جدول ۲۱ عملکرد مدل‌ها را با استفاده از معیارهای کارایی ذکر شده بیان می‌کند. با مقایسه نتایج، معلوم شد که مدل ACBR با بردار تعبیه roberta دارای بالاترین دقت طبقه‌بندی برای هر دو مجموعه داده پایه و مجموعه داده شامل شکلک است که دقت آن

مقایسه مدل ACBR با نتایج بررسی شده قبلی در جدول ۲۲ نشان داده شده است.

مراجع

- [1]. S. Rosenthal, N. Farra, P. Nakov, "Sentiment analysis in twitter", In Proceedings of the 9th international workshop on semantic evaluation (SemEval 2015), pp. 451- 463, Dec 2019.
- [2]. S Poria, E. Cambria, N. Howard, GB. Huang and A. Hussain, "Fusing audio, visual and textual clues for sentiment analysis from multimodal content", Neurocomputing, Volume 174, Part A, Pages 50-59, 22 January 2016.
- [3]. P. Chikersal, S. Poria, E. Cambria, A. Gelbukh and C.E. Siong, "Modelling public sentiment in twitter: using linguistic patterns to enhance supervised learning", in: International Conference on Intelligent Text Processing and Computational Linguistics, LNTCS, Springer, vol. 9042, pp. 49-65, 2015.
- [4]. S. Kiritchenko, X. Zhu and SM. Mohammad, "Sentiment analysis of short informal texts". Journal of Artificial Intelligence Research, vol. 50, pp. 723-762, Aug 20, 2014.
- [5]. Y. Kim, "Convolutional neural networks for sentence classification", In Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP). pp. 1746-1751, 25 Aug 2014.
- [6]. E. Cambria, S. Poria, A. Gelbukh and M. Thelwall, "Sentiment analysis is a big suitcase". IEEE Intelligent Systems, vol. 32, pp. 74-80, December 2017.
- [7]. Y. Zhang and B. Wallace, "A sensitivity analysis of (and practitioners' guide to) convolutional neural networks for sentence classification" IJCNLP, 13 Oct 2015.
- [8]. C. dos Santos, M. Gatti, "Deep convolutional neural networks for sentiment analysis of short texts". In Proceedings of COLING 2014, the 25th International Conference on Computational Linguistics, pp. 69-78, August 2014.
- [9]. M. Cliché, "Twitter sentiment analysis with cnns and lstms", Published in Proceedings of SemEval-2017, arXiv:1704.06125, 8 pages, 20 Apr 2017.
- [10]. T. Mikolov, K. Chen, G. Corrado and J. Dean, "Efficient Estimation of word representations in vector space" arXiv:1301.3781, Jan 2013.
- [11]. A. Hogenboom, D. Bal, F. Frasincar, M. Bal, F. de Jong and U. Kaymak, "Exploiting emoticons in sentiment analysis", in: Proceedings of the 28th Annual ACM Symposium on Applied Computing, pp. 703-710, March 2013.
- [12]. N. Hu, P.A. Pavlou, and J. Zhang, "Can online reviews reveal a product's true quality? Empirical findings and analytical modeling of online Word-of-Mouth communication", EC '06: Proceedings of the 7th ACM conference on Electronic commerce, Pages 324-330, June 2006.
- [13]. M. Pontiki, D. Galanis, H. Papageorgiou, I. Androutsopoulos, S. Manandhar and M. Mahmoud Al-Ayyoub, "SemEval-2016 Task 5: Aspect Based Sentiment Analysis", in: 10th International Workshop

جدول ۲۲ - مقایسه مدل های تحلیل احساسات

نویسنده / سال	مدل	مجموعه داده	دقت (%)
بصیری و همکاران (۲۰۲۰)	ABCDM	Sentiment140	۸۱/۸۲
مورنو و همکاران (۲۰۲۱)	L-C	Sentiment140	۸۴/۲
کامیاب و همکاران (۲۰۲۱)	ACL-SA	Sentiment140	۸۷/۱۲
آیتوگ اونان (۲۰۲۲)	RCNNGWE	Sentiment140	۸۴/۰۶
مدل پیشنهادی	ACBR	Sentiment140	۹۰/۱۸

۷- نتیجه گیری و پیشنهاد برای تحقیقات آتی

این مقاله نظرات نوشته شده به زبان انگلیسی که حاوی شکلک بودند را در نظر گرفت. برای پرداختن به این موضوع، از مجموعه داده sentiment140 استفاده شد و شکلک ها برای تعیین قطبیت احساسات در دو کلاس مثبت و منفی در نظر گرفته شدند. مجموعه داده توپتر با استفاده از دو آزمایش موردسجش قرار گرفت، که هر دو مجموعه داده بدست آمده با مدل های شبکه عصبی CNN و LSTM آزمایش شدند. در این مدل ها از چهار بردار تعبیه از قبیل آموزش دیده word2vec، GloVe، BERT و roberta استفاده شد. در نهایت، عملکرد مدل ها با استفاده از چند معیار عملکرد مختلف، مانند دقت، صحت، پوشش، امتیاز F1 و میانگین کلان مقایسه شدند. پس از بررسی معیارهای عملکرد برای همه مدل ها، نتایج نشان داد که مدل ACBR با بردار تعبیه از پیش آموزش دیده roberta، دقت بهتری را هنگام طبقه بندی احساسات برای این دو مجموعه داده ارائه می کند. همچنین نتایج نشان داد که شکلک ها ارزش بیشتری را هنگام تعیین طبقه بندی احساسات اضافه می کنند. از آنجایی که کاربران بیشتری هستند که نظرات و بازخورد خود را با استفاده از شکلک ها ارائه می دهند، باید این سیگنال ها را هنگام شناسایی احساسات در نظر گرفت.

کارهای آینده شامل موارد زیر است:

- (۱) در نظر گرفتن بازخورد یا نظرات چندزبانه کاربران
- (۲) توسعه مدل هایی که می توانند طعنه را تشخیص دهند
- (۳) در نظر گرفتن تصاویر، صدا و کلیپ های ویدئویی، علاوه بر متن برای طبقه بندی احساسات

- analysis", *Future Generation Computer Systems*, Vol. 115, pp. 279-294 September 2020.
- [27]. M. Kamyab, G. Liu and M. Adjeisah, "Attention-Based CNN and Bi-LSTM Model Based on TF-IDF and GloVe Word Embedding for Sentiment Analysis", *Applied Sciences*, Vol. 11, November 2021.
- [28]. A. Onan, "Bidirectional convolutional recurrent neural network architecture with group-wise enhancement mechanism for text sentiment classification", *Journal of King Saud University - Computer and Information Sciences*, Vol 34, pp. 2098-2117, May 2022.
- [29]. M.A. Ullah, S.M. Marium, Sh. Begum, N.S. Dipa, "An algorithm and method for sentiment analysis using the text and emoticon", *ICT Express*, Vol. 6, pp. 606–615, December 2020.
- [30]. B.Pang, , L.Lillian, and S. Vaithyanathan, "Thumbs Up? Sentiment classification using machine learning techniques", In *Proceedings of Conference on Empirical Methods in Natural Language Processing (EMNLP)*, *Computation and Language*, pp. 79–86, May 2002.
- [31]. M. Hu, and B. Liu, "Mining and summarizing customer reviews" In *Proceedings of ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD)*, pp. 168-177, August 2004.
- [32]. D. Sharma, M. Sabharwal, V. Goyal and M. Vij, "Sentiment analysis techniques for social media data: A review", *First International Conference on Sustainable Technologies for Computational Intelligence*, Springer, pp.75-90, January 2020.
- [33]. k. jagmeet and M. sabharwal, "Spam detection in online social networks using feed forward neural network", In: *RSRI Conference on Recent Trends in Science and Engineering*, pp. 69–78, 2018.
- [34]. M. Unnisa, A. Ameen and S. Raziuddin. "Opinion mining on twitter data using unsupervised learning technique", *International Journal of Computer Applications*, vol 148(12), 2016.
- [35]. X. Zhang, J. Zhao and Y. LeCun. "Character-level convolutional networks for text classification", In *Advances in Neural Information Processing Systems 28*, *Neural Information Processing Systems Foundation*. pp. 649–657, 2015.
- [36]. A. Karpathy, J. Johnson and L. Fei-Fei. "Visualizing and understanding recurrent networks", *ICLR*, arXiv:1506.02078, Jun 2015.
- [37]. C. Baziotis, N. Pelekis and C. Doulkeridis, "Datastories at semeval 2017 task 4: Deep lstm with attention for message-level and topic-based sentiment analysis", In *Proceedings of the 11th International Workshop on Semantic Evaluation*, Association for Computational Linguistics, pp. 747-754, August 2017.
- [38]. J. Devlin, M.W. Chang, K. Lee and K. Toutanova, "Bert: Pre-training of deep bidirectional transformers for language understanding", *Proceedings of NAACL-HLT*, pp: 4171–4186, Oct 2018.
- [39]. A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A.N. Gomez, L. Kaiser and I. Polosukhin, "Attention is all you need," in *Proc. 31st Conference on Neural Information Processing Systems (NIPS)*, Long Beach, pp. 6000–6010, 2017.
- on *Semantic Evaluation (SemEval)*, pp.19 – 30, Jan 2019
- [14]. Y. Wan and Q. Gao, "An ensemble sentiment classification system of twitter data for airline services analysis", *IEEE International Conference on Data Mining Workshop (ICDMW)*, Atlantic City, NJ, USA, pp. 1318–1325, November 2015.
- [15]. A. Pak and P. Paroubek. "Twitter as a corpus for sentiment analysis and opinion mining", *Conference: Proceedings of the International Conference on Language Resources and Evaluation, LREC 2010*, vol. 13, pp. 1320–1326, 17-23 May 2010.
- [16]. W. Wolny, "Emotion analysis of twitter data that use emoticons and emoji ideograms", *The 25th International Conference on Information Systems Development (ISD2016)*, *European Language Resources Association (ELRA)*, Valletta, Malta, Corpus ID: 31794814, 2016.
- [17]. E. Kouloumpis, T. Wilson and J. Moore, "Twitter sentiment analysis: The good the bad and the omg!", in: *Fifth International AAAI Conference on Weblogs and Social Media*, Vol. 5, pp. 538-541, 2011.
- [18]. D. Jurafsky and J. H. Martin, "Speech and language processing: An introduction to natural language processing, computational linguistics, and speech recognition", In *Prentice Hall series in artificial intelligence*, Prentice Hall, 2000.
- [19]. J. Lee, W. Yoon, S. Kim, D. Kim, C. H. So and J. Kang, "Biobert: A pre-trained biomedical language representation model for biomedical text mining", *Bioinformatics*, Oxford, England, vol. 36(4), pp. 1234–1240, 15 February 2020.
- [20]. Y. Goldberg and O. Levy, "word2vec Explained: deriving Mikolov et al.'s negative-sampling word-embedding method", arXiv:1402.3722, 15 Feb 2014.
- [21]. X. Fu, W. Liu, Y. Xu and L. Cui, "Combine HowNet lexicon to train phrase recursive autoencoder for sentence-level sentiment analysis". *Neurocomputing*, Vol 241, pp. 18–27, 7 June 2017.
- [22]. E. Bender and A. Lascarides, "Linguistic fundamentals for natural language processing II: 100 essentials from semantics and pragmatics". *Synthesis Lectures on Human Language Technologies*, vol. 12(3), 1–268 pages, November 2019.
- [23]. Y. Sharma, G. Agrawal, P. Jain and T. Kumar, "Vector representation of words for sentiment analysis using GloVe", *International Conference on Intelligent Communication and Computational Techniques (ICCT)*, IEEE, ISBN:978-1-5386-3030-3, December 2017.
- [24]. A. Varshney, Y. Kapoor, A. Thukral, R. Sharma and P. Bedi, "Performing Sentiment Analysis on Twitter Data Using Deep Learning Models: A Comparative Study" *Advances in Data and Information Sciences*, vol 318, pp 371–381, 08 February 2022.
- [25]. C.N. Dang, M.N. Moreno-García and F. Prieta, "Hybrid Deep Learning Models for Sentiment Analysis", Vol. 2021, 16 pages, Aug 2021.
- [26]. M.E. Basiri, Sh. Nemati , M. Abdar, E. Cambria and U.R. Acharya, "ABCDM: An Attention-based Bidirectional CNN-RNN Deep Model for sentiment

[40]. Attention layer Retrieved from:

<https://github.com/deepak-kaji/mimic-lstm>

زیر نویس ها

25. Accuracy
26. Precision
27. Recall
28. HTML
29. Natural Language Toolkit
30. Tokenizer
31. Devlin
32. Next Sentence Prediction(NSP)
33. Self attention
34. Embedding layer
35. Convolutional layer
36. regularization
37. Dense layer
38. Cliche
39. Attention Baseb Cnn And Lstm Model With RoBERT

1. Sentiment Analysis
2. Emoticons
3. Convolutional neural network (CNN)
4. Long short-term memory (LSTM)
5. Word embedding
6. Natural language processing
7. Context
8. Document level
9. Support vector machine(SVM)
10. Naive bayes
11. One-hot
12. Term frequency-inverse document frequency(TF-IDF)
13. Bidirectional long short-term memory
14. Bidirectional encoder representations from transformers (BERT)
15. Sentence level
16. Aspect level
17. Data preprocessing
18. features selection
19. Stemming
20. Stop words
21. Supervised learning
22. Unsupervised learning
23. Anisa
24. maximum entropy