



بهبود نظر کاوی فارسی مبتنی بر قطبیت و

متوازن سازی کلمات مهم مثبت و منفی (مطالعه

موردی: نظرات دیجی کالا برای موبایل)

مهديه واحدي پور، محبوبه شمسي و عبدالرضا رسولی كناري

دانشكده مهندسي برق و كامپيوتر، دانشگاه صنعتی قم، قم - ایران

چکیده

بسیاری از شبکه‌های اجتماعی و سایت‌ها به مردم اجازه می‌دهند تا احساسات و نظرات خود را در مورد محصولات و خدمات مختلف به اشتراک بگذارند. در این مقاله روشی جدید مبتنی بر قطبیت نظرات مثبت و منفی فارسی درباره محصولات تلفن همراه از سایت دیجی کالا و داده‌های سنتی پرس ارائه شده است. نتیجه اجرا با الگوریتم‌های بیز ساده، ماشین بردار پشتیبان، کاهش گرادیان تصادفی، رگرسیون لجستیک، جنگل تصادفی و یادگیری عمیق مانند شبکه عصبی کانولوشن و حافظه کوتاه مدت متوالی بر اساس پارامترهایی مانند صحت، بازیابی، معیار فیشر و دقت، مورد توجه قرار گرفته شده است. روش پیشنهادی روی داده‌های دیجی کالا، با الگوریتم‌های بیز ساده بین ۱۰ تا ۳۴ درصد و ماشین بردار پشتیبان بین ۵ تا ۲۴ درصد و کاهش گرادیان تصادفی بین ۷ تا ۳۸ درصد و رگرسیون لجستیک بین ۵ تا ۳۸ درصد و جنگل تصادفی بین ۴ تا ۲۲ درصد و روش شبکه عصبی کانولوشن به میزان ۴ درصد افزایش دقت را به همراه داشته است. همچنین در داده‌های سنتی پرس با الگوریتم‌های بیز ساده بین ۱۲ تا ۴۶ درصد و ماشین بردار پشتیبان بین ۵ تا ۴۶ درصد و کاهش گرادیان تصادفی بین ۵ تا ۳۵ درصد و رگرسیون لجستیک بین ۶ تا ۴۶ درصد و جنگل تصادفی بین ۴ تا ۴۶ درصد دقت نسبت به قبل از اعمال روش پیشنهادی به دست آمده است.

واژگان کلیدی: تحلیل احساسات، نظر کاوی، یادگیری ماشین، یادگیری عمیق، قطبیت

Improving Persian Opinion Mining based on Polarity and Balancing Positive and Negative Keywords (case study: Digikala reviews for mobile)

Mahdieh Vahedipoor¹, Mahboubeh Shamsi and Abdolreza Rasouli Kenari

Faculty of Electrical and Computer Engineering, Qom University of Technology, Qom, Iran

Abstract:

In recent years, the massive growth of generated content by the users in social networks and online marketing sites, allows people to share their feelings and opinions in a variety of opinions about different products and services. Sentiment analysis is an important factor for better decision-making that is done using natural language processing (NLP), computational methods, and text analysis to extract the polarity of unstructured documents. The complexity of human languages and sentiment analysis have created a challenging research context in computer science and computational linguistics. Many researchers used supervised machine learning algorithms such as Naïve Bayes (NB), Stochastic Gradient Descent (SGD), Support Vector Machine (SVM), Logistic Regression (LR) Random Forest (RF), and deep learning algorithms such as Convolution Neural Network (CNN) and Long Short-Term Memory (LSTM). Some researchers have used Dictionary-based methods. Despite the existence of effective techniques in text mining, there are still unresolved challenges. Note that user comments are

* Corresponding author

* نویسنده عهده‌دار مکاتبات

unstructured texts; Therefore, in order to structure the textual inputs, parsing is usually done along with adding some features, linguistic interpretations and removing additional items, and inserting the next terms in the database, then extracting the patterns in the structured data and finally the outputs will evaluate and interpret. The imbalance of data with the difference in the number of samples in each class of a dataset is an important challenge in the learning phase. This phenomenon breaks the performance of the classifications because the machine does not learn the features of the unpopulated classes well. In this paper, words are weighted based on the prescribed dictionary to influence the most important words on the result of the opinion mining by giving higher weight. On the other hand, the combination of the adjacent words using n-gram methods will improve the outcome. The dictionaries are highly related to the domain of the application. Some words in an application are important but in mobile comments are not impressive. Another challenge is the unbalanced train data, in which the number of positive sentences is not equal to the number of negative sentences. In this paper, two ideas are applied to build an efficient opinion mining algorithm. First, we build a precise dictionary for mobile Persian comments, and the second idea is to balance the positive and negative comments in train data. In summary, the main achievements of the current research can be mentioned: creating a weighted comprehensive dictionary in the field of mobile phone opinions to increase the accuracy of opinion analysis, balancing positive and negative opinions to improve the accuracy of opinion analysis, and eliminating the negative effect of overfitting and providing a precise approach to Determining the polarity of users' opinions about mobile phones using machine learning and recurrent deep learning algorithms. This new method is presented on mobile phone products from the Digikala site and Senti-Pers data. The result is performed with Naive Bayesian, Support Vector Machine, Stochastic Gradient Descent, Logistic Regression, Random Forest, and deep learning methods such as Convolutional Neural Network and Long Short-Term Memory based on parameters such as Accuracy, Precision, Retrieval, and F-Measure. The proposed method increases accuracy on Digikala, with NB between 10% and 34% and SVM between 5% and 24%, SGD between 7% and 38%, LR between 5% to 38%, and RF between 4% Up to 22% and CNN by 4%. The results show an accuracy increment on Senti-Pers, with NB between 12% and 46% and SVM between 5% and 46%, SGD between 5% and 35%, LR between 6% to 46%, and RF between 4% Up to 46%.

Keywords: Sentiment Analysis, Opinion Mining, Machine Learning, Deep Learning, Polarity

به‌طور عمده دو نوع تکنیک یادگیری ماشین برای دسته‌بندی متون وجود دارد: تکنیک مبتنی بر یادگیری ماشین تحت نظارت و تکنیک مبتنی بر یادگیری ماشین بدون نظارت. در روش یادگیری نظارت‌شده، مجموعه‌داده‌ها توسط کارشناسان به‌صورت دستی یا با استفاده از فرهنگ لغت‌نامه برچسب‌گذاری می‌شوند [۵،۲۳]. بنابراین یادگیری مدل‌ها با بخشی از داده‌ها به‌عنوان داده‌های آموزشی صورت گرفته و با بخش دیگری از داده‌ها به‌عنوان داده‌های آزمایشی یک مدل ریاضی معقول را به دست می‌آورند [۲]. برخلاف یادگیری نظارت‌شده، در فرآیند یادگیری بدون نظارت هیچ اطلاعاتی در مورد برچسب وجود ندارد و نمی‌توان به‌راحتی پردازش کرد، زیرا آموزشی صورت نمی‌گیرد. الگوریتم‌های خوشه‌بندی با دسته‌بندی داده‌های مشابه جهت حل مشکل پردازش داده‌های بدون برچسب به کار می‌روند [۶].

از سوی دیگر الگوریتم‌های یادگیری عمیق، زیرمجموعه‌ای از الگوریتم‌های یادگیری ماشین هستند که هدف آن‌ها کشف چندین سطح از نمایش‌های توزیع‌شده از داده ورودی است. در سال‌های اخیر الگوریتم‌های یادگیری عمیق زیادی برای حل مسائل هوش مصنوعی سنتی ارائه شده‌اند. کاربردهای آن‌ها در زمینه‌های مختلف

۱- مقدمه

در حال حاضر شبکه‌ها و رسانه‌های اجتماعی نقش مهمی در ارائه اطلاعات در مورد هر محصول با استفاده از نظرات مختلف کاربران دارند. در علوم رایانه، تجزیه و تحلیل احساسات به‌عنوان نظرکاوی شناخته شده است و نظرات مردم و همچنین احساسات آن‌ها نسبت به محصولات، سازمان‌ها و ویژگی‌های مربوط به آن‌ها را تحلیل می‌کند. محققان با روش‌های مختلف یادگیری ماشین و به‌منظور استخراج اطلاعات معنادار از احساسات مردم از این داده‌ها استفاده می‌کنند. تشخیص ساختار متن به سه شکل مختلف سند، جمله و کلمه امکان‌پذیر است. در سطح سند، اسناد به نظرات مثبت (موافق)، منفی (مخالف) و خنثی (بی‌طرف) دسته‌بندی می‌شود. در سطح جمله تعیین می‌کنند که آیا جمله نظر مثبت، منفی و یا خنثی را بیان می‌کند و در سطح ویژگی، برای تعیین نظر مثبت، منفی و یا خنثی به جزئی‌ترین شکل ممکن (کلمه) انجام می‌گیرد و نظرات را در مورد ویژگی‌های خاص موضوع بررسی می‌کنند [۱]. در این مقاله، داده‌های جمع‌آوری شده نظرات تلفن همراه از سایت دیجی‌کالا و داده‌های سنتی پرس از سایت پیکره‌گان تحلیل احساسات در سطح جمله موردتوجه قرار گرفته است.

بینایی ماشین همانند دسته‌بندی تصاویر، شناسایی اشیاء، استخراج تصاویر، قطعه‌بندی معنایی و برآورد ژست انسان است [۳].

صرف‌نظر از الگوریتم مورد استفاده برای یادگیری، نظرات کاربران، متون بدون ساختار هستند؛ بنابراین، برای ساختاردهی به ورودی‌های متنی به طور معمول تجزیه، همراه با افزودن برخی ویژگی‌ها، تفاسیر زبانی و حذف موارد اضافی و درج موارد بعدی در پایگاه داده انجام می‌گیرد، سپس استخراج الگوهای درون داده‌های ساختاریافته و درنهایت ارزیابی و تفسیر خروجی‌ها انجام می‌شود. عدم تعادل داده‌ها با اختلاف در تعداد نمونه‌ها در هر کلاس از یک مجموعه داده چالش مهمی در مرحله یادگیری است. این پدیده عملکرد دسته‌بندی‌ها را خراب می‌کند، زیرا ماشین ویژگی‌های کلاس‌های با جمعیت کمتر را خوب نمی‌آموزد. استفاده از تکنیک‌های نمونه‌گیری از جمله نمونه‌گیری بیش‌ازحد و نمونه‌گیری کمتر از حد عملکرد دسته‌بندی را بهبود می‌بخشد. برای مجموعه داده‌های کوچک، نمونه‌گیری بیش‌ازحد انتخاب شده است که مناسب‌ترین استراتژی است به دلیل این‌که مجموعه اولیه نمونه را افزایش می‌دهد و این کار، بر روی وظایف دودویی یا چند ستونی تمرکز دارد [۴]. در این مقاله از نمونه‌گیری بیش‌ازحد برای سازگاری داده‌ها استفاده شده است.

در این پژوهش، ابتدا کلمات توقف و سایر اطلاعات ناخواسته از نظرات کاربرانی که با استفاده از خزنده‌ها جمع‌آوری شده‌اند، حذف می‌گردد و سپس طبق فرهنگ‌نامه مورد نظر ارزش نظرات، محاسبه شده و قطبیت جدیدی به نظرات داده می‌شود. به دلیل ناهمگن بودن داده‌های مثبت و منفی با روش نمونه‌گیری بیش‌ازحد، داده‌ها متعادل می‌شوند و از طریق فرآیند بردار سازی داده‌های متن ساختاریافته استخراجی از مرحله پیش‌پردازش به ماتریس اعداد تبدیل می‌شوند؛ سپس این ماتریس‌های عددی به الگوریتم‌های مختلف یادگیری ماشین تحت نظارت و یادگیری عمیق برای دسته‌بندی این نظرات داده شده و سپس پارامترهای مختلف مانند صحت، بازیابی، معیار فیشر و دقت برای ارزیابی عملکرد الگوریتم‌های یادگیری ماشین تحت نظارت و یادگیری عمیق جهت بهبود نظرکاوی نظرات فارسی مورد مقایسه قرار گرفته است. به‌طور خلاصه از دستاوردهای اصلی تحقیق حاضر می‌توان به موارد زیر اشاره کرد:

- ساخت یک فرهنگ‌نامه جامع وزن‌دار در حوزه نظرات تلفن همراه جهت افزایش دقت نظرکاوی
 - متوازن سازی نظرات مثبت و منفی جهت بهبود دقت نظرکاوی و از بین بردن اثر منفی بیش‌برازش
 - ارائه رویکرد دقیق جهت تعیین قطبیت نظرات کاربران در مورد تلفن همراه با استفاده از الگوریتم‌های یادگیری ماشین و یادگیری عمیق بازگشتی
- ساختار مقاله به شرح زیر است: در بخش ۲، به بررسی ادبیات تحقیق پرداخته شده است. در بخش ۳، پیاده‌سازی رویکرد پیشنهادی توضیح داده شده است. در بخش ۴، ارزیابی عملکرد روش پیشنهادی انجام شده است و دقت رویکرد پیشنهادی با کارهای دیگران مقایسه شده است و در انتها در بخش ۵ مقاله نتیجه‌گیری و کارهای آینده پیشنهاد شده است.

۲- کارهای گذشته

تجزیه و تحلیل احساسات در حوزه وسیعی مانند بررسی فیلم، بررسی تدریس، بررسی محصولات، یادگیری الکترونیکی، بررسی هتل و بسیاری دیگر مورد مطالعه قرار گرفته است. بیشتر دانشمندان متمرکز به تجزیه و تحلیل داده‌های کمی هستند [۲۱-۲۵]. با این حال، برخی از مطالعات انجام‌شده بر روی داده‌های کیفی با استفاده از تجزیه و تحلیل احساسات انجام‌شده است و همچنین مطالعه خوبی توسط نویسندگان مختلف بر اساس دسته‌بندی احساسات در سطح سند انجام شده است که در زیر به برخی از آن‌ها اشاره شده است. Lee و Pang دسته‌بندی احساسات در سطح سند با احساسات مثبت و منفی را تحقیق و بررسی کرده‌اند. آن‌ها با سه الگوریتم مختلف یادگیری ماشین نظیر SVM، NB و ME، آزمایش کرده و فرآیند دسته‌بندی را با استفاده از تکنیک‌های n-gram مانند تک‌واژه‌ای، دوواژه‌ای و ترکیب تک‌واژه‌ای و دوواژه‌ای انجام دادند. آن‌ها چارچوب ویژگی‌های کیفی را برای اجرای الگوریتم‌های یادگیری ماشین استفاده کرده‌اند. در نتیجه تجزیه و تحلیل آن‌ها، الگوریتم NB نتایج ضعیف و الگوریتم SVM نتایج بهتری را نشان می‌دهد [۷]. Salvetti و همکارانش، در مفهوم کلی قطبیت عقیده‌ای (OvOp) با استفاده از الگوریتم‌های یادگیری ماشین نظیر مدل بیز ساده و مدل مارکوف را برای دسته‌بندی متون مورد بحث قرار داده‌اند و آن‌ها توسط WordNet و POS متون را برچسب‌گذاری کرده‌اند و

به عنوان فیلتر لغوی برای دسته بندی عمل می کنند. آزمایش آن ها نشان می دهد که نتیجه به دست آمده از فیلتر WordNet نسبت به فیلتر POS کمتر است. رویکرد آن ها، نتیجه بهتری روی داده های وب نشان می دهد [۸]. Beineke و همکارانش از مدل بیز ساده برای دسته بندی احساسات استفاده کرده اند. آن ها ویژگی های مشتق شده را برای پیش بینی احساسات استخراج کرده اند. آن ها ویژگی های مشتق شده دیگری را به منظور بهبود نتیجه دقت به مدل اضافه کرده اند و داده های مورد نظر را برای تعیین تأثیر نسبی استفاده می کنند. این ایده به سیستم اجازه می دهد که به عنوان یک مدل احتمالی در مقیاس منطقی خطی عمل کند. در این روش هر واژه الحاقی با ضریب k همراه است که این ضرایب باید شناخته شده باشند. با این حال، زمانی که اسناد برچسب دار در دسترس هستند، ممکن است تخمین آن ها مفید باشد [۹]. Mullen و Collier با استفاده از الگوریتم ماشین بردار پشتیبان ارزش ها را به کلمات انتخاب شده اختصاص می دهند تا یک مدل برای دسته بندی ایجاد کنند. علاوه بر این، کلاس های مختلفی از ویژگی های نزدیک به موضوع اختصاص داده شده با مقادیر بهتر به کار برده اند که به دسته بندی کمک می کند. نویسندگان مقایسه ای از رویکرد پیشنهادی خود را با داده ها، حاشیه نویسی موضوع و حاشیه نویسی دستی ارائه کرده اند. رویکرد پیشنهادی در مقایسه با حاشیه نویسی موضوع نشان داده است که در مقایسه با داده های حاشیه نویسی دستی نتایج بهبود بیشتری دارد. یک مشکل این روش با محدود کردن دامنه با اضافه کردن محدودیت های کلمه مرتبط به موضوع این است که تعداد شمارش نتیجه به شدت کاهش می یابد و هرگونه افزایش بالقوه را از بین می برد. هم چنین به نظر می رسد که جنبه های موضوع روابط تحقیق حاضر فقط سطحی بیان می شود [۱۰]. Dave و همکارانش تکنیک های بازایی اطلاعات مورد استفاده برای بازایی ویژگی و نتیجه روش های مختلف را مورد آزمایش قرار داده اند. آن ها از یک ابزار برای سنتز نظرات استفاده کرده اند، سپس آن ها را تغییر داده و با استفاده از سایت های تجمعی دسته بندی می کنند. آن ها با استفاده از این نظرات، ویژگی ها را شناسایی کرده و در نهایت روش هایی برای تعیین مثبت یا منفی بودن نظرات به کار می برند و برای دسته بندی جملات حاصل از جستجوی وب با استفاده از نام محصول به عنوان شرایط جستجو استفاده می کنند. در این روش مواردی چون عدم تطابق

درجه بندی، هم بستگی، مقایسه را دشوار می سازد [۱۱]. Matsumoto و همکارانش از ارتباطات نحوی میان کلمات به عنوان مبنای تحلیل احساسات در سطح جمله استفاده کرده اند. در پژوهش خود، توالی کلمه مکرر و وابستگی جملات را با استفاده از الگوریتم SVM استخراج می کنند. آن ها روش های تک واژه ای، دو واژه ای، توالی کلمه و وابستگی زیردرخت از هر جمله در مجموعه داده ها را استخراج می کنند. این روش در سطح سند امکان پذیر نیست [۱۲]. Su و Xu و Zhang روشی برای به دست آوردن ویژگی های معنایی با استفاده از word2vec برای گرفتن ویژگی های مشابه و سپس دسته بندی نظرات از طریق الگوریتم SVM ارائه داده اند. رویکرد آن ها بر اساس دو بخش است. در بخش اول، آن ها از ابزار word2vec برای خوشه بندی ویژگی های مشابه استفاده کرده اند تا ویژگی های معناداری در دامنه انتخاب شده را ذخیره کنند. سپس در بخش دوم، با انتخاب ویژگی مبتنی بر واژگان به ساخت داده های آموزشی می پردازند. این روش روی زبان چینی انجام شده است [۱۳]. Chen و Liu چندین دسته بند مختلف را در تحلیل احساسات پیشنهاد کرده اند. آن ها از یازده روش دسته بندی چندسطحی در دو مجموعه داده های میکرو بلاگ و هشت ماتریس ارزیابی مختلف برای تجزیه و تحلیل استفاده کرده اند. آن ها هم چنین از سه فرهنگ لغت مختلف احساسات برای دسته بندی چندسطحی استفاده می کنند. این روش از فرهنگ نامه از پیش تعریف شده استفاده کرده است در حالی که می توان در دامنه های مختلف فرهنگ نامه مورد نظر را ایجاد کرد و به دلیل برچسب های چندگانه مقایسه آزمایش های مختلف گران است [۱۴]. Luo و همکارانش روشی برای تبدیل متن به فضای عاطفی کمتر (ESM)، با استفاده از تکنیک های یادگیری ماشین جهت دسته بندی متون را پیشنهاد کرده اند. آن ها متن را به کلماتی که معنی قطعی و روشن دارند تبدیل می کنند و برای دسته بندی کلمات آن را به شش دسته اصلی مانند خشم، ترس، انزجار، غم و اندوه، شادی و تعجب تقسیم می کنند. آن ها دو روش متفاوت برای تخصیص وزن به کلمات را با برچسب های عاطفی در نظر گرفته اند و تمام وزن کلمات عاطفی را محاسبه کرده اند و بر اساس این مقادیر، پیام ها به گروه های مختلف دسته بندی می شوند. این روش می تواند نتایج منطقی را در هر مجموعه داده یا دامنه استفاده کند ولی تبدیل کلمات به معنای قطعی و تقسیم آن ها به شش دسته اصلی زمان بر است [۱۵]. Niu

و همکارانش مجموعه‌ای از داده‌ها را که حاوی مجموعه‌ای از جفت‌متن تصویر با حاشیه‌نویسی دستی از تویتر است جمع‌آوری کرده‌اند. رویکرد آن‌ها از تجزیه و تحلیل احساسات دارای دو بخش یادگیری مبتنی بر واژگان و آماری است که در مورد تجزیه و تحلیل مبتنی بر واژگان، مجموعه‌ای از کلمات یا عبارات در نظر گرفته شده است و نمره احساسات از پیش تعریف شده دارند. درحالی که در یادگیری آماری، تکنیک‌های مختلف یادگیری ماشین با ویژگی‌های متن مورد استفاده قرار می‌گیرد [۱۶]. Junming و Caiqiang مدل ارزیابی معلم آموزشی آنلاین را بر اساس نظرکاوی انجام می‌دهند. آن‌ها با استفاده از خزنده وب، نظرات دانش‌آموزان را که در زبان چینی منتشر شده و در سیستم مدیریت یادگیری نوشته شده است، جمع‌آوری می‌کنند. اگر کلمات احساسی متن ذهنی در لغت‌نامه قطب‌واژه نیستند، آن‌ها از روش متقابل اطلاعات متناوب برای قضاوت قطبیت استفاده کردند. این مدل یک ارزیابی کلی از هر معلم ارائه می‌دهد [۱۷]. Martin, Ortigosa و Carro روش دیگر تحلیل احساسات برای محیط یادگیری الکترونیکی را با استفاده از روش ترکیبی از تکنیک‌های یادگیری زبان اسپانیایی پیشنهاد دادند. در آزمایش آن‌ها ترکیبی از تکنیک‌های مبتنی بر واژگان و SVM بالاترین دقت را به دست می‌آورند. در زمینه یادگیری الکترونیکی، ممکن است اطلاعاتی در مورد احساسات دانش‌آموز از پیام‌هایی که در فیس‌بوک می‌نویسند، استخراج شود. احساسات دانش‌آموزان نسبت به یک دوره می‌تواند به عنوان بازخورد برای معلمان، به‌ویژه در مورد یادگیری آنلاین، مفید باشد. با این حال، این کار همچنان دارای محدودیتی برای تحلیل احساسات است، تمام کلمات به عنوان یک قطب، نظرات را برچسب‌گذاری می‌کنند، تمام کلمات مثبت نمره ۱، تمام کلمات منفی نمره -۱ و تمام کلمات خنثی نمره صفر را می‌گیرند آن‌ها وزن‌های مختلف را به کلمات مختلف اختصاص نمی‌دهند و برای بدتر و بد امتیاز مشابه به دست می‌آورند [۱۸]. Rungworawut و Pong-inwong پیشنهاد ساخت تدریس تئوری ارزیابی احساسات را برای قطبیت احساسات خودکار ارائه دادند. کار آن‌ها متشکل از دو بخش است، بخش اول تهیه داده‌ها و بخش دوم مدل‌سازی داده‌ها و ارزیابی است. نویسندگان انتخاب ID3، بیزین و ماشین بردار پشتیبان را برای دسته‌بندی ارزیابی احساسات آموزشی انتخاب کردند. نتایج تجربی آن‌ها نشان می‌دهد که SVM دارای بالاترین دقت است. در این کار، نمره وزن

توسط یک متخصص تعریف شده است که دارای تجربه در ارزیابی تدریس است. وزن احساسات از ۱۰۰- به ۱۰۰+ متغیر بود. روش پیشنهادی آن‌ها نمی‌تواند واژه تقویت‌کننده را در نظر بگیرد و در زبان تایلندی ساخته شده است [۱۹]. Hee Yong Youn و Yili Wang روش وزن‌دهی تطبیقی جدید به نام میان‌رده قدرت تشخیصی با مدل Chi-square پیشنهاد کرده‌اند که مقیاس صحیح میان‌رده و اهمیت داخل دسته‌ای را مشخص می‌کند. یک مدل ریاضی جدید برای اندازه‌گیری قدرت تشخیصی میان‌رده‌ها از ویژگی‌ها ارائه کرده‌اند، درحالی که یک مدل Chi-square اصلاح‌شده برای اندازه‌گیری وابستگی ویژگی‌ها در داخل دسته‌بندی ارائه شده است. همچنین استراتژی خوشه‌بندی ریزنمونه برای تعیین حاکمیت ویژگی‌های تشخیصی و تعیین ویژگی‌های توزیع مشابه برای وزن‌دهی کارایی پیشنهاد شده است. علاوه بر این، یک استراتژی وزن سازگاری پیشنهاد شده است تا وزن هر ویژگی را به درستی تعیین کند. به‌طور خاص، توزیع یک کلمه در یک دسته خاص نشان‌دهنده وابستگی آن به آن دسته است که با توجه به وابستگی وزن، دقیق‌تر می‌تواند باشد. از آنجاکه طرح پیشنهادی بر اساس وابستگی هر یک از ویژگی به رده مربوطه است، طول مدت جمله به‌طور مستقیم بر عملکرد روش پیشنهادی تأثیر نمی‌گذارد و همان‌طور که هر کلمه دارای وزن مشخصی است منعکس‌کننده درجه اهمیت آن برای دسته‌های مختلف است، قطبیت یک جمله دارای دو دیدگاه مختلف است که توسط وزن‌های به دست آمده از دسته‌بندی متن تعیین می‌شود. برای عقاید مختلف یا متضاد ایجاد نشده است. در طرح پیشنهادی شکلک‌ها با استفاده از تابعی در نرم‌افزار متلب شناسایی و حذف می‌شوند درحالی که احساسات به‌طور معمول احساسات مردم را در برمی‌گیرند و این روش مناسب نیست [۲۰].

۳- رویکرد پیشنهادی

داده‌کاوی یکی از گرایش‌های بسیار مهم در علم کامپیوتر است. در داده‌کاوی، سعی بر این است که دانش از حجم انبوهی از داده‌ها استخراج شود. منظور از دانش، اطلاعات مطلوبی است که مفید و هدفمند و از قبل شناخته شده‌اند. در این میان نظرکاوی رشته‌ای است که به‌تازگی بسیار مورد توجه قرار گرفته است و تحقیقات زیادی برای بررسی و توسعه تکنیک‌های مربوط به نظرکاوی انجام شده است. نظرکاوی قانونی در تقاطع بازیابی اطلاعات و محاسبات

زبان‌شناسی است که نظرات کاربران در مورد یک موضوع داده‌شده را کاوش می‌کند. هدف نظرکاوی این است که نظر نویسندگان و یا یک نظر کلی در مورد یک سند خاص را استخراج کند و راهکارهایی برای دسته‌بندی نظرات مندرج در سایت‌ها و وبلاگ‌ها است. در حقیقت در نظرکاوی به دنبال یافتن تابعی هستیم که بتواند با دریافت ورودی یک کامنت آن را از منظر احساسات به یکی از مقادیر مثبت و منفی نگاشت نماید. این مسئله در معادله ۱ فرموله شده است.

$$Label = f_{opinion-mining}(Comment) \quad (1)$$

$$Label \in \{Pos, Neg\}$$

که در آن *Comment* نظر کاربر و *Label* نتیجه نظرکاوی شامل یکی از مقادیر مثبت و منفی خواهد بود. از آنجاکه برآورد همچنین تابعی از نظر روش‌های مرسوم ریاضی امکان‌پذیر نخواهد بود، لذا استفاده از الگوریتم‌های یادگیری ماشین جهت تقریب تابع فوق مفید خواهد بود. چنانچه بخواهیم از الگوریتم‌های یادگیری ماشین برای این منظور استفاده کنیم نیاز به یک مجموعه داده آموزشی خواهیم داشت که برچسب گذاری شده باشد. مراحل رویکرد پیشنهادی در (شکل ۱) مورد بحث قرار گرفته است.



(شکل - ۱). نمایشی از روش پیشنهادی
(Figure - 1). Proposed Method Schema

ابتدا با استفاده از یک خزنده وب طراحی شده در زبان پایتون نظرات فارسی کاربران در مورد انواع تلفن همراه استخراج و در یک مجموعه نظرات غیر ساخت‌یافته ذخیره می‌گردد. مجموعه نظرات غیر ساخت‌یافته استخراج شده در معادله ۲ فرموله شده است.

$$UnCmts = \{Cmt_1, Cmt_2, \dots, Cmt_n\} \quad (2)$$

که در آن Cmt_i نظر کاربر i -ام و n تعداد کل نظرات استخراج شده است.

در مرحله پیش‌پردازش، برای بهبود تعیین قطبیت (مثبت و منفی) یک نظر، داده‌های اضافی از نظرات، مانند کلمات پایانی و کلمات غیر کلیدی مانند از، اگر، به و غیره، اعداد، علائم نقطه‌گذاری، حروف انگلیسی، حروف و علائم غیرفارسی، حذف فاصله‌ها، حروف تکراری بیش از دو بار

در یک کلمه، کلماتی با طول یک که معنی خاصی ندارند و نمی‌توان برای آن‌ها قطبیت تعیین کرد را از متن اصلی برای بهبود نظرکاوی و تعیین دقیق قطبیت حذف کرده‌ایم که این عملیات را پیش‌پردازش متن می‌نامیم.

به‌طور خلاصه با توجه به الگوریتم (۱) ارائه‌شده برای پیش‌پردازش متون نظرات تلفن همراه عملیات زیر انجام شده است:

- حذف اعداد و علائم و کلمات غیر حروف فارسی مانند علائم نقطه‌گذاری، تعجب، اصطلاحات انگلیسی و غیره (خط ۲، ۳ و ۴).
- حذف حروف تکراری بیش از یک‌بار در کلمه به‌طور مثال به جای کلماتی چون عالللللی کلمه عالی جایگزین می‌شود (خط ۵).
- حذف کلمات یک‌حرفی که در جمله بی‌معنا هستند (خط ۶).
- نرمال‌سازی متن که به‌منظور حذف فاصله‌ها است (خط ۷).
- جداسازی کلمات به‌منظور استخراج ویژگی‌ها استفاده می‌شود (خط ۸).
- حذف کلمات پایانی مانند اگر، به، چون و غیره که آن‌ها در تعیین احساسات نقش مهمی ایفا نمی‌کنند (خط ۹).

ALGORITHM Preprocessing Pseudo Code

INPUT: Unstructured Comments (*UnCmts*)

Dataset of the User's Comments

OUTPUT: Structural Comments (*C*)

Structured Dataset of the User's Comments

1. FOR EACH c IN *UnCmts*
2. Remove Digits
3. Remove Punctuation Mark
4. Remove Non-Persian letter
5. Remove Repeat Characters
6. Remove Repeat Words
7. Remove Spaces (Normalization)
8. Tokenize Words
9. Remove Persian StopWords
10. Add c To Structured Comments *C*
11. Return *C*

(الگوریتم-۱). پیش پردازش

مجموعه اعمال مربوط به پیش‌پردازش نظرات در معادله ۳ فرموله شده است.

$$C = Preprocess(UnCmts)$$

$$C = \{C_1, C_2, \dots, C_n\} \quad (3)$$

$$C_i = \{kw_{1i}, kw_{2i}, \dots, kw_{m_i}\}$$

که در آن C مجموعه نظرات ساختاریافته و n تعداد نظرات و kw_i کلمات کلیدی حذف نشده از نظر i -ام و m_i تعداد کلمات کلیدی نظر i -ام هستند. در تجزیه و تحلیل

که در آن *Dic* فرهنگ‌نامه‌ای شامل اجتماع کلیه لغات کلیدی در نظرات کاربران است و تابع *Score* به ازای هر لغت کلیدی عددی بین ۵+ و ۵- را نگاشت می‌کند. نبود یک فرهنگ‌نامه وزن‌دهی به کلمات در زبان فارسی از لحاظ بار احساسی به‌ویژه در زمینه تلفن همراه از چالش‌های مهم این پژوهش به شمار می‌آید. این فرهنگ‌نامه بیشتر لغات استفاده‌شده توسط کاربران در نظرسنجی تلفن همراه را جمع‌آوری و از لحاظ بار احساسی امتیازدهی نموده است.

مرحله بعدی تعیین قطبیت هر یک از نظرات کاربران و برچسب گذاری آن‌ها جهت ایجاد داده آموزشی است. این کار با استفاده از فرهنگ‌نامه ساخته‌شده در مرحله قبل صورت می‌پذیرد. الگوریتم (۲) جزییات برچسب‌گذاری را بیان می‌کند.

ALGORITHM Labeling Pseudo Code

INPUT: Structural Comments (*C*)
Structured Dataset of the User's Comments
Keyword's Dictionary (*Dic*)
Score function

OUTPUT: Labeled User's Comments (*L*)

1. FOR EACH *C* in *C*
2. FOR EACH *kw* in *C*
3. *S* = Sum all *Score(kw)*
4. IF *S* > 0 THEN *L* = POS
5. IF *S* < 0 THEN *L* = NEG
6. ELSE remove *C* from *C*
7. Add (*C*, *L*) to *L*
8. Return *L*

(الگوریتم-۲). برچسب گذاری

در این مرحله طبق الگوریتم برچسب‌گذاری، هر نظر را به‌طور جداگانه (خط ۱) به‌صورت کلمه به کلمه تفکیک کرده (خط ۲) و امتیاز هر کلمه را از فرهنگ‌نامه ساخته‌شده در مرحله قبل (خط ۳) پیدا می‌کنیم. اگر کلمات موجود در نظر در فرهنگ‌نامه وجود داشت، امتیاز آن کلمات را با هم جمع می‌کنیم (خط ۳). اگر مجموع امتیازات یک نظر بیشتر از صفر باشد آن نظر برچسب مثبت (خط ۴) و اگر کمتر از صفر باشد برچسب منفی را به خود اختصاص می‌دهد (خط ۵)؛ بنابراین کل مجموعه داده به این روش برچسب‌گذاری می‌شوند.

روال برچسب‌گذاری در معادله ۵ فرموله شده است. *L* مجموعه زوج مرتب‌های (نظر، برچسب) است به‌طوری‌که چنانچه مجموع امتیازات کلمات کلیدی هر نظر بزرگ‌تر از صفر شود برچسب مثبت و اگر کوچک‌تر از صفر شود برچسب منفی خواهد خورد. همچنین m_C تعداد کلمات کلیدی هر نظر است.

$$L = \{(C, L) \mid C \in C, L \in \{POS, NEG\}\} \quad (5)$$

احساسات، مدل چندواژه‌ای مانند تک‌واژه‌ای، دوواژه‌ای، سه‌واژه‌ای و برای نمونه‌های بزرگ‌تر چهارواژه‌ای و پنج-واژه‌ای به تجزیه و تحلیل احساسات متن یا سند کمک می‌کند. به‌عنوان مثال روش‌های چندواژه‌ای را با استفاده از جمله "The movie is not a good one" شرح می‌دهیم. در روش تک‌واژه‌ای جمله فوق به‌صورت تک کلمه به شرح 'The' و 'movie' و 'is' و 'not' و 'a' و 'good' و 'one' است. در روش دوواژه‌ای جفت کلمات در نظر گرفته شده است. بنابراین کلمات کلیدی جمله فوق در این روش به شرح 'The movie is' و 'movie is' و 'is not' و 'not a' و 'a good one' و 'good' است. در روش سه‌واژه‌ای مجموعه‌ای از کلمات سه‌تایی در نظر گرفته می‌شود. بنابراین کلمات کلیدی جمله فوق در این روش به شرح 'The movie is a good one' و 'movie is not a good one' و 'is not a good one' و 'a good one' است.

بیشتر نویسندگان از روش تک‌واژه‌ای برای طبقه‌بندی نظرات استفاده کرده‌اند. این رویکرد نتیجه‌ای بهتر را فراهم می‌کند چون وابستگی به کلمات قبل خود ندارد، اما در بعضی موارد شکست می‌خورد به‌عنوان مثال در جمله "The movie is not a good one" هنگامی که با استفاده از روش تک‌واژه‌ای تجزیه و تحلیل می‌شود، قطبیت جمله باوجود یک قطب مثبت کلمه "good" و یک قطب منفی کلمه "not" خنثی می‌شود؛ اما هنگامی که این جمله با استفاده از رویکرد دوواژه‌ای تجزیه و تحلیل می‌شود، قطبیت جمله به دلیل وجود کلمات "not good" منفی تشخیص داده می‌شود؛ بنابراین، انتظار می‌رود که نتیجه بهتر باشد. پس از پیش‌پردازش نظرات، از کلمات کلیدی باقی‌مانده در کل داده‌ها، یک فرهنگ‌نامه به‌صورت دستی ایجاد می‌کنیم و به هر یک از کلمات با توجه به کاربرد آن‌ها امتیازی بین ۵- تا ۵+ می‌دهیم. درواقع کلمات خیلی ضعیف منفی با نمره (۵-)، کلمات ضعیف منفی با نمره (۴-)، خصوصیات منفی تلفن همراه با نمره (۳-)، کلمات متوسط منفی با نمره (۲-)، کلمات تقریباً نادرست منفی با نمره (۱-)، کلمات نامفهوم با نمره (۰)، کلمات تقریباً نادرست مثبت با نمره (۱)، کلمات متوسط مثبت با نمره (۲)، خصوصیات مثبت تلفن همراه با نمره (۳)، کلمات قوی مثبت با نمره (۴) و کلمات خیلی قوی مثبت با نمره (۵) امتیازدهی می‌شوند. ساخت فرهنگ‌نامه در معادله ۴ فرموله شده است.

$$\begin{aligned} Dic &= \bigcup_{i=1}^n C_i = \{all \text{ unique } kw \text{ in } C\} \\ Score: Dic &\rightarrow \{-5, -4, \dots, +4, +5\} \\ \text{for example} \\ Score(عالی) &= +5, Score(افتضاح) = -5 \end{aligned} \quad (4)$$

الگوریتم (۴) بردارسازی ابتدا برای هر نظر مقدار فراوانی کلمات کلیدی آن محاسبه می‌شود. سپس این مقدار با تقسیم بر تعداد کل فراوانی نرمال می‌گردد (خط ۱). مقدار اهمیت هر کلمه در کل نظرات نیز با محاسبه لگاریتم معکوس فراوانی آن کلمه در کل نظرات تقسیم بر فراوانی در آن نظر به دست می‌آید (خط ۲). ضرب این دو مقدار میزان اهمیت یک کلمه در یک نظر را مشخص می‌کند. به این ترتیب کلماتی که در بیشتر نظرات وجود دارند اهمیت تصمیم‌گیری کمتری پیدا می‌کنند. در انتها به ازای همه نظرات و همه کلمات موجود در فرهنگ‌نامه مقدار اهمیت در ماتریس T ذخیره می‌گردد (خط ۳-۶). این ماتریس عددی به عنوان داده آموزشی و آزمون برای دسته‌بندها مورد استفاده قرار خواهد گرفت.

ALGORITHM Vectorizer Pseudo Code

INPUT: Balanced Labeled User's Comments (L^*)
Structural Comments (C)
Keyword's Dictionary (Dic)
OUTPUT: Labeled TF/IDF Matrix (T)

1. $TF(kw, C)$: Return $\left(\frac{freq(kw)}{\text{number of words in } C}\right)$
2. $IDF(kw, C)$:
Return $\log\left(\frac{\text{number of comments } |C|}{\text{number of comments including } kw}\right)$
3. FOR EACH (C, L) in L^*
4. FOR EACH kw in Dic
5. $T[C, kw] = TF(kw, C) \times IDF(kw, C)$
6. $T[C, class] = L$
7. Return T

(الگوریتم-۴). بردارسازی

مراحل ساخت ماتریس در معادله ۷ فرموله شده است؛ که در آن $freq$ نمایانگر فراوانی یک کلمه در سند است. بقیه متغیرها مانند معادلات قبل هستند.

$$TF(kw, C) = \left(\frac{freq(kw)}{\sum_{kw^* \in C} freq(kw^*)}\right)$$

$$IDF(kw, C) = \log\left(\frac{|C|}{|\{C \in C: kw \in C\}|}\right)$$

$$T_{n \times (m+1)}: \text{Matrix of } (n = |L^*|)$$

$$\begin{aligned} & \times (m+1) \\ & = |Dic| + 1 \text{ elements} \end{aligned} \quad (V)$$

$$\forall (L_i, C_i) \in L^* \quad i = 1..n$$

$$\forall kw_j \in Dic \quad j = 1..m :$$

$$T[i, j] = TF(kw_j, C_i) \times IDF(kw_j, C)$$

$$T[i, m+1] = L_i$$

پس از ساخت ماتریس داده‌ها نوبت به دسته‌بندی داده‌ها با الگوریتم‌های یادگیری ماشین است الگوریتم (۵). به این منظور داده‌ها به دو بخش آموزشی (۷۰٪) و آزمون (۳۰٪) به صورت تصادفی تقسیم می‌شوند (خط ۲). این ماتریس‌ها طبق الگوریتم دسته‌بندی با پنج الگوریتم مختلف یادگیری ماشین تحت نظارت از جمله بیز ساده (خط ۳)، ماشین بردار پشتیبان (خط ۴)، کاهش گرادیان تصادفی (خط ۵)، رگرسیون لجستیک (خط ۶) و جنگل

$$L = \begin{cases} POS & \sum_{i=1}^{m_c} Score(kw_i) > 0 \\ NEG & \sum_{i=1}^{m_c} Score(kw_i) < 0 \end{cases}$$

چالش مهم بعدی در نظرات استخراج‌شده از سایت‌ها، عدم توازن بین نظرات مثبت و منفی است که می‌تواند بر نتیجه الگوریتم یادگیری تأثیر منفی داشته باشد. به دلایل مختلف مثل بازاریابی، تبلیغات، مسائل فرهنگی، وجود الفاظ نامناسب و ... بسیاری از نظرات منفی از سایت‌ها حذف می‌شوند و در نتیجه تعداد نظرات منفی استخراج‌شده عموماً کمتر از نظرات مثبت است. از این رو از روش‌های متعادل‌سازی، از جمله روش نمونه‌گیری بیش‌ازحد و روش نمونه‌گیری کمتر از حد برای بهبود عملکرد دسته‌بندی می‌توان استفاده نمود. در این پژوهش از روش نمونه‌گیری بیش‌ازحد استفاده شده است زیرا این روش برای مجموعه داده‌های کوچک مناسب‌تر است و مناسب‌ترین استراتژی جهت افزایش مجموعه نمونه اولیه است و این کار را با تمرکز بر روی نمونه‌های دودویی یا چند ستونی انجام می‌دهد. در این روش نمونه‌های بیشتری به «کلاس کم‌جمعیت‌تر» اضافه می‌شوند. البته در تولید نمونه‌های جدید سعی می‌شود بیشترین شباهت با نمونه‌های موجود صورت پذیرد تا توزیع آماری نمونه‌ها تغییر نکند. با توجه به الگوریتم ۳ متوازن‌سازی نمونه‌های جدید با ترکیب تصادفی نمونه‌های موجود در داده‌ها با برچسب منفی ساخته می‌شوند (خط ۳-۵). این کار آن‌قدر ادامه پیدا می‌کند تا تعداد نمونه‌های هر دو کلاس نسبتاً برابر شوند (خط ۶).

ALGORITHM Balancing Pseudo Code

INPUT: Labeled User's Comments (L)
OUTPUT: Balanced Labeled User's Comments (L^*)

1. $L^* = L$
2. REPEAT
3. Choose two random C_1, C_2 from L where $L = NEG$
4. $C^* = \text{Combination of } C_1, C_2$
5. Add (C^*, NEG) to L^*
6. UNTIL $||[POS]|| = ||[NEG]||$
7. Return L^*

(الگوریتم-۳). متوازن‌سازی

روال متوازن‌سازی در معادله ۶ فرموله شده است.

$$L^* = L \cup \{(C^*, NEG)\} \quad (۶)$$

$$C^*: \text{Mixture of 2 random neg comments}$$

در این مرحله برای استفاده از الگوریتم‌های یادگیری ماشین و یادگیری عمیق برای ارزیابی عملکرد روش پیشنهادی داده‌های متنی باید به بردارهای عددی تبدیل شوند. در این تحقیق با استفاده از روش بردار ویژگی که اهمیت کلمه را به سند نشان می‌دهد برای ایجاد بردار عددی استفاده شده است. بردارهای عددی حاصل از مجموعه نظرات تشکیل ماتریس T را می‌دهند. طبق

در بخش بعدی به ارزیابی رویکرد پیشنهادی روی دو مجموعه داده نظرات جمع آوری شده از سایت دیجی کالا در مورد تلفن همراه و داده های سنتی پرس می پردازیم.

۴- ارزیابی عملکرد رویکرد پیشنهادی

در این بخش ابتدا مجموعه داده های استفاده شده، سپس معیارهای ارزیابی و در نهایت سناریوهای مختلف ارزیابی رویکرد پیشنهادی مورد بررسی قرار گرفته است. کلیه ارزیابی ها توسط زبان برنامه نویسی پایتون انجام شده است [۲۶].

۴-۱- مجموعه داده

در این پژوهش، مجموعه داده های فارسی نظرات تلفن همراه از سایت دیجی کالا و مجموعه داده های سنتی پرس از سایت پیکره گان برای تحلیل احساسات در نظر گرفته شده است. این مجموعه داده دیجی کالا را با استفاده از خزنده وب جمع آوری کرده ایم که شامل ۳۴۶۲۵ نظرسنجی با برچسب مثبت و ۵۰۸۵ نظرسنجی با برچسب منفی است. همچنین مجموعه داده سنتی پرس از سایت پیکره گان قابل دریافت است که شامل ۷۴۱۰ نظرسنجی با برچسب مثبت و ۲۴۹۵ نظرسنجی با برچسب منفی و ۳۴۱۵ نظرسنجی با برچسب خنثی است. در این پژوهش برای بهبودی نظرکاوی یک فرهنگ نامه از کلمات مثبت و منفی داده های نظرات تلفن همراه ایجاد کرده ایم و به هر یک از کلمات با توجه به ارزش آن ها وزنی اختصاص داده ایم. این فرهنگ نامه برای داده های دیجی کالا شامل ۱۴۰۹۰ لغت با وزن های مختلف از ۵- تا ۵+ است که دارای ۷۲۷۰ کلمه با خصوصیت مثبت و ۶۸۲۰ کلمه با خصوصیت منفی است. مجموعه داده های دیجی کالا بعد از به دست آوردن قطبیت جدید نظرات طبق فرهنگ نامه ارائه شده دارای ۳۱۸۹۵ نظرسنجی با برچسب مثبت و ۶۶۱۵ نظرسنجی با برچسب منفی است. همچنین فرهنگ نامه برای داده های سنتی پرس شامل ۱۴۲۰۵ لغت با وزن های مختلف از ۵- تا ۵+ است که دارای ۳۳۱۵ کلمه با خصوصیت مثبت و ۱۳۱۰ کلمه با خصوصیت منفی و ۹۵۸۰ کلمه نامفهوم است. مجموعه داده های سنتی پرس بعد از به دست آوردن قطبیت جدید نظرات طبق فرهنگ نامه ارائه شده دارای ۷۴۱۰ نظرسنجی با برچسب مثبت و ۲۴۹۵ نظرسنجی با برچسب منفی است. سپس به دلیل ناهمگن بودن داده های مثبت و منفی با استفاده از روش نمونه گیری بیش از حد، کل داده ها را به بالاترین مقدار سوق می دهیم و داده های

تصادفی (خط ۷) و همچنین از الگوریتم های یادگیری عمیق مانند شبکه عصبی کانولوشن (خط ۸) و حافظه کوتاه مدت متوالی (خط ۹) یادگیری و مقادیر دقت، صحت، بازیابی و معیار فیشر هر کدام ذخیره می گردد. برای دستیابی به دقت بالاتر و از بین بردن اثر بیش برزش این کار به تعداد دفعات مختلف اجرا و میانگین نتایج جهت ارزیابی مورد استفاده قرار می گیرد (خط ۱۰).

ALGORITHM Classification Pseudo Code

INPUT: Labeled TF/IDF Matrix ($\mathcal{T}$)

OUTPUT: Accuracy, Precision, Recall, F-Measure
DeepAccuracy, DeepLoss

1. Repeat NumFolds
2. [Train, Test] = Random Split of $\mathcal{T}$ (70%, 30%)
3. [Accuracy, Precision, Recall, F-Measure] = NB(Train, Test)
4. [Accuracy, Precision, Recall, F-Measure] = SVM(Train, Test)
5. [Accuracy, Precision, Recall, F-Measure] = SGD(Train, Test)
6. [Accuracy, Precision, Recall, F-Measure] = LR(Train, Test)
7. [Accuracy, Precision, Recall, F-Measure] = RF(Train, Test)
8. [DeepAccuracy, DeepLoss] = CNN($\mathcal{T}$)
9. [DeepAccuracy, DeepLoss] = LSTM($\mathcal{T}$)
10. [Accuracy, Precision, Recall, F-Measure] / NumFolds
11. Return [Accuracy, Precision, Recall, F-Measure]
12. Return [DeepAccuracy, DeepLoss]

الگوریتم-۵). دسته بندی

مفروضات زیر جهت کلیه دسته بندی ها برقرار است.

Train the Classifier using Training Data $\mathcal{T}$

$\mathcal{T} = \{\mathcal{T}_i\}_{i=1..m}$, $\mathcal{T}_i = (\mathbf{T}_i, L_i)$, $\mathbf{T}_i = [TF.IDF(kw_j)]_{j=1..m}$

$L_i \in \{pos, neg\}$

Classifier $F: \mathbf{T} \rightarrow L$

که در آن بردار $\mathbf{T}$ به عنوان ورودی دسته بند شامل مقادیر TF.IDF کلمات کلیدی یک نظر و برچسب L به عنوان خروجی تابع دسته بندی در نظر گرفته شده است. به ازای نظرات جدید (دیده نشده) دسته بند آموزش دیده سعی در پیش بینی برچسب می کند.

for a given comment (Cmt):

$C = Preprocessing(Cmt)$ s.t. $C = \{kw_j\}_{j=1..m_C}$

$\mathbf{T} = Vectorizer(C) = [TF.IDF(kw_j)]_{j=1..m}$

predict $\hat{L} = F_{\theta}(\mathbf{T})$

Error = $|L - \hat{L}|$

بدین ترتیب که ابتدا نظر پس از پیش پردازش به مجموعه ای از کلمات کلیدی تبدیل و سپس مقادیر TF.IDF کلمات محاسبه و در بردار $\mathbf{T}$ ذخیره می شود. مقدار قطبیت نظر توسط دسته بند آموزش دیده شده محاسبه می گردد ($\hat{L}$). این مقدار بعداً جهت ارزیابی دقت و صحت الگوریتم قابل استفاده است.

سیستم را به سمت صحت یا بازیابی بهینه‌سازی کنیم که تأثیر بیشتری بر نتیجه نهایی دارد. [۳۶]

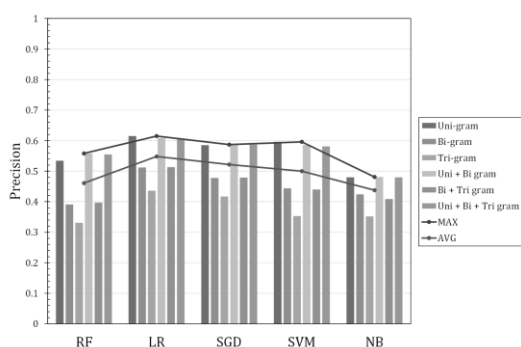
(جدول ۱-). فراوانی نتایج صحیح و اشتباه دسته‌بندی
(Table - 1). True/False Classification result frequency

برچسب‌گذاری دسته‌بندی	مثبت	منفی
مثبت	مثبت واقعی (TP)	مثبت کاذب (FP)
منفی	منفی واقعی (TN)	منفی کاذب (FN)
دقت	$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}$	
صحت	$Precision = \frac{TP}{TP + FP}$	
بازیابی	$Recall = \frac{TP}{TP + FN}$	
معیار فیشر	$F - Measure = \frac{2 * Precision * Recall}{Precision + Recall}$	

۳-۴- تأثیر استفاده از فرهنگ‌نامه و

متوازن‌سازی بر داده‌های دیجی‌کالا

در این قسمت نتایج پارامترهای ارزیابی نظرات تلفن همراه از سایت دیجی‌کالا با استفاده از الگوریتم‌های یادگیری ماشین توسط تکنیک‌های تک‌واژه‌ای، دوواژه‌ای، سه‌واژه‌ای، ترکیب تک‌واژه‌ای و دوواژه‌ای، ترکیب دوواژه‌ای و سه‌واژه‌ای و ترکیب هر سه تکنیک بررسی شده است. با استفاده از ماتریس درهم‌ریختگی، پارامترهای مختلف ارزیابی مانند دقت، صحت، بازیابی و معیار فیشر به‌دست‌آمده پس از اعمال روش پیشنهادی محاسبه شده است. دسته‌بندی نتایج بر اساس انواع تکنیک‌های چندواژه‌ای و با استفاده از الگوریتم‌های مختلف یادگیری ماشین در Error! Reference source not found. (۲) نشان داده شده است.



دیجی‌کالا با برچسب مثبت و منفی هر یک دارای ۳۱۸۹۵ و داده‌های سنتی‌پرس با برچسب مثبت و منفی هر یک دارای ۷۴۱۰ نظرسنجی می‌شوند. بعد از آماده‌سازی داده‌ها از ۷۰ درصد آن‌ها برای آموزش و ۳۰ درصد برای آزمایش استفاده کرده‌ایم. داده‌های آزمایش از روش ۱۰ دور استفاده می‌کند یعنی داده‌های آزمایش در ۱۰ جای مختلف از کل داده آزمایش می‌شود و میانگین دقت را برمی‌گرداند.

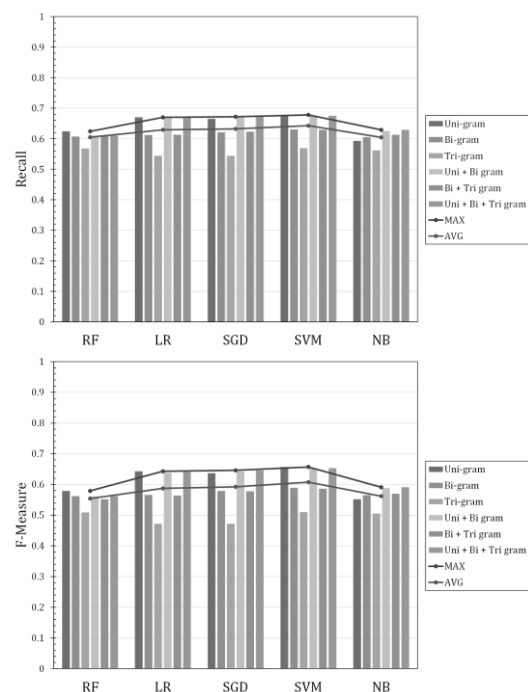
۲-۴- پارامترهای ارزیابی عملکرد

پارامترهای مفید برای ارزیابی عملکرد الگوریتم یادگیری ماشین تحت نظارت از یک ماتریس شناخته‌شده به‌عنوان ماتریس درهم‌ریختگی یا جدول احتمالی استخراج می‌شوند. این پارامترها در الگوریتم یادگیری ماشین تحت نظارت مورد استفاده قرار می‌گیرند که در ارزیابی عملکرد هر الگوریتم کمک می‌کند. از نظر دسته‌بندی، برای مقایسه برچسب کلاس‌ها اصطلاحات مثبت واقعی (TP)، مثبت کاذب (FP)، منفی واقعی (TN)، منفی کاذب (FN) محاسبه می‌شوند. مثبت واقعی، تعداد نظراتی را نشان می‌دهد که مثبت هستند و همچنین در دسته‌بندی به‌صورت مثبت دسته‌بندی می‌شوند، مثبت کاذب نشان‌دهنده نظرات مثبت است اما دسته‌بندی آن را به‌صورت مثبت دسته‌بندی نمی‌کند. به‌طور مشابه، منفی واقعی نشان‌دهنده نظرات منفی است که توسط دسته‌بندی منفی دسته‌بندی شده‌اند و منفی کاذب، نظرات منفی هستند اما دسته‌بندی آن به‌صورت منفی دسته‌بندی نمی‌شود [۳۶]. این مقادیر در Error! Reference source not found. (۱) نشان داده شده است.

بر اساس مقادیر به‌دست‌آمده از جدول فراوانی، برای ارزیابی عملکرد هر دسته‌بندی، پارامترهای دقت، صحت، بازیابی و معیار فیشر به دست می‌آیند. دقت رایج‌ترین پارامتر فرآیند دسته‌بندی است که می‌توان به‌عنوان نسبت نمونه‌های دسته‌بندی‌شده مثبت به تعداد کل نمونه‌ها محاسبه شود. صحت دقیق بودن نتیجه دسته‌بندی را اندازه‌گیری می‌کند و نسبت تعداد نمونه‌هایی که به‌طور صحیح، مثبت برچسب‌گذاری شده‌اند به تعداد کل نمونه‌هایی که مثبت دسته‌بندی شده‌اند. بازیابی کامل بودن نتیجه دسته‌بندی را اندازه‌گیری می‌کند و نسبت تعداد کل نمونه‌های مثبت برچسب‌گذاری شده به نمونه‌های کاملی است که واقعاً مثبت هستند. معیار فیشر هارمونیک متوسط صحت و بازیابی است. لازم است که

ترکیبی تکواژه‌ای و دوواژه‌ای به دلیل وجود روش تک-واژه‌ای و دوواژه‌ای در نتیجه نسبت به روش‌های دیگر نیز نتایج بهتری را به دست می‌آورند. مقدار دقت به‌دست‌آمده با تکنیک‌های دوواژه‌ای و سه‌واژه‌ای در الگوریتم ماشین بردار پشتیبان بین ۱ تا ۳ درصد بهتر از مقدار حاصل‌شده در الگوریتم‌های دیگر است. روش ترکیبی دوواژه‌ای و سه‌واژه‌ای به دلیل تأثیر روش سه‌واژه‌ای باعث می‌شود که مقدار دقت نسبت به روش‌های دیگر نیز نتایج ضعیف‌تری را به دست آورد. در نهایت مقدار دقت به‌دست‌آمده با سه تکنیک در الگوریتم‌های ماشین بردار پشتیبان و کاهش گرادیان تصادفی و رگرسیون لجستیک بین ۶ تا ۹ درصد بهتر از مقدار حاصل‌شده در الگوریتم‌های دیگر هستند. روش ترکیبی تکواژه‌ای و دوواژه‌ای و سه‌واژه‌ای به دلیل تأثیر روش سه‌واژه‌ای باعث می‌شود که مقدار دقت نسبت به روش‌های دیگر نیز نتایج ضعیف‌تر و تأثیر روش تک-واژه‌ای و دوواژه‌ای نتایج بهتری را به دست آورند؛ بنابراین نتایج خوبی را به ارمغان می‌آورند.

به‌طور خلاصه مقدار دقت به‌دست‌آمده در الگوریتم‌های ماشین بردار پشتیبان در همه روش‌ها به‌غیر از روش سه‌واژه‌ای به میزان ۷ درصد و کاهش گرادیان تصادفی و رگرسیون لجستیک در همه روش‌ها به‌غیر از دوواژه‌ای و سه‌واژه‌ای و روش ترکیبی دوواژه‌ای و سه‌واژه‌ای به میزان ۵ درصد بهتر از مقدار حاصل‌شده در الگوریتم‌های دیگر هستند. تأثیر روش سه‌واژه‌ای باعث می‌شود که مقدار دقت نسبت به روش‌های دیگر در الگوریتم‌های مختلف نیز نتایج ضعیف‌تری را به دست‌آورند. از سوی دیگر صحت دقیق بودن نتیجه دسته‌بندی است بنابراین هر چه مقدار صحت بیشتر باشد عملکرد دسته‌بندی دقیق‌تر خواهد بود. همان‌طور که مشاهده می‌شود مقدار صحت به‌دست‌آمده در الگوریتم رگرسیون لجستیک در همه روش‌ها به‌غیر از روش دوواژه‌ای و سه‌واژه‌ای و ترکیب دوواژه‌ای و سه‌واژه‌ای به میزان ۲ درصد بهتر از مقدار حاصل‌شده در الگوریتم‌های دیگر هستند. همچنین بازیابی کامل بودن نتیجه دسته‌بندی است بنابراین نسبت به صحت هر چه مقدار بازیابی کمتر باشد عملکرد دسته‌بندی کامل‌تر خواهد بود. مقدار بازیابی به‌دست‌آمده روش سه‌واژه‌ای در همه الگوریتم‌های یادگیری ماشین به میزان ۱۱ درصد نتایج بهتری از مقدار حاصل‌شده در روش‌های دیگر را به دست می‌آورد. معیار فیشر متوسط صحت و بازیابی است. لازم است که سیستم را به سمت صحت یا بازیابی بهینه‌سازی کنیم که تأثیر

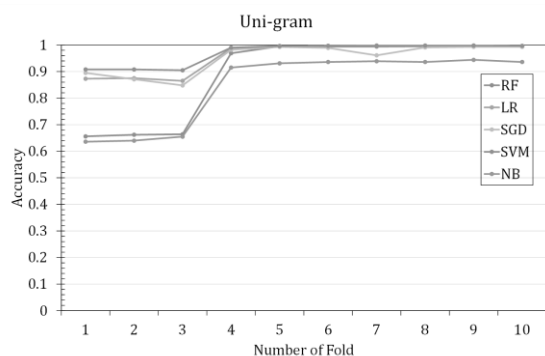


(شکل ۲). نتایج پارامترهای ارزیابی الگوریتم‌های یادگیری ماشین بر اساس تکنیک‌های چندواژه‌ای (دیجی کالا)
(Figure - 2). Machine learning algorithms result based on n-gram technics on Digikala comments

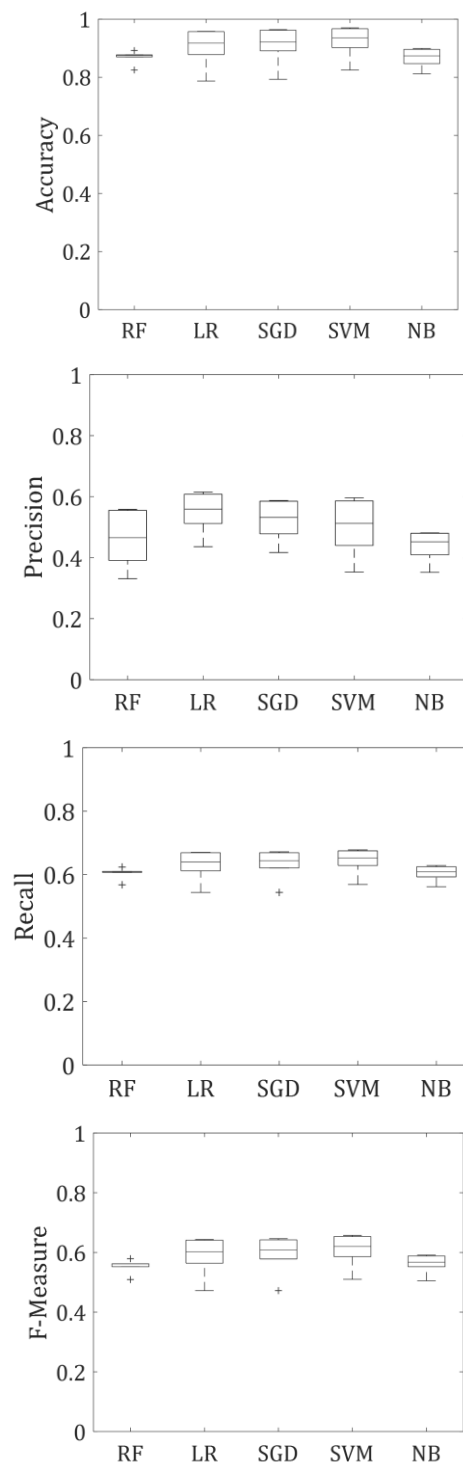
همان‌طور که در Error! Reference source not found. (۲) نشان داده شده است، مقدار دقت به‌دست‌آمده با تکنیک تکواژه‌ای در الگوریتم ماشین بردار پشتیبان بین ۲ تا ۱۳ درصد بهتر از مقدار حاصل‌شده در الگوریتم‌های دیگر است. روش تکواژه‌ای به دلیل این‌که به‌صورت کلمه به کلمه پردازش می‌شود در نتیجه نسبت به روش‌های دیگر نتایج بهتری را به دست می‌آورد. مقدار دقت به‌دست‌آمده با تکنیک دوواژه‌ای در الگوریتم ماشین بردار پشتیبان بین ۱ تا ۳ درصد بهتر از مقدار حاصل‌شده در الگوریتم‌های دیگر است. روش دوواژه‌ای به دلیل این‌که به‌صورت دوکلمه‌ای پردازش می‌شود در نتیجه نسبت به روش‌های دیگر به‌غیر از تکواژه‌ای نتایج بهتری را به دست می‌آورد. همچنین مقدار دقت به‌دست‌آمده با تکنیک سه‌واژه‌ای در الگوریتم‌های ماشین بردار پشتیبان و جنگل تصادفی بین ۱ تا ۴ درصد بهتر از مقدار حاصل‌شده در الگوریتم‌های دیگر هستند. روش سه‌واژه‌ای به دلیل این‌که به‌صورت سه‌کلمه‌ای پردازش می‌شود و کلمات چند بار تکرار می‌شوند؛ بنابراین، بر احتمال سند اثر می‌گذارد در نتیجه نسبت به روش‌های دیگر دقت دسته‌بندی را کاهش می‌دهد. مقدار دقت به‌دست‌آمده با تکنیک‌های تکواژه‌ای و دوواژه‌ای در الگوریتم‌های ماشین بردار پشتیبان و کاهش گرادیان تصادفی بین ۱ تا ۹ درصد بهتر از مقدار حاصل‌شده در الگوریتم‌های دیگر هستند. روش

(Figure - 3). Statistical charts of different algorithms such as Average, SD, and Quarters (Digikala)

دقت به‌دست‌آمده پس از اعمال روش پیشنهادی و با استفاده از الگوریتم‌های مختلف یادگیری ماشین با آزمایش در ۱۰ دور متفاوت در Error! Reference source not found. (۴) نشان داده شده است. مشاهده می‌شود که مقدار دقت به‌دست‌آمده در الگوریتم‌های ماشین بردار پشتیبان به میزان ۲۵ درصد و کاهش گرادیان تصادفی به میزان ۲۴ درصد و رگرسیون لجستیک به میزان ۲۲ درصد از دور اول و الگوریتم‌های بیز ساده به میزان ۲۶ درصد و جنگل تصادفی به میزان ۳۰ درصد از دور چهارم به سمت بالا بهبود یافته‌اند. دقت به‌دست‌آمده در روش دوواژه‌ای پس از اعمال روش پیشنهادی و با استفاده از الگوریتم‌های مختلف یادگیری ماشین در همه الگوریتم‌های یادگیری ماشین از دور چهارم بین ۲۳ تا ۳۴ درصد به سمت بالا بهبود یافته‌اند. دقت به‌دست‌آمده در روش سه‌واژه‌ای پس از اعمال روش پیشنهادی از دور چهارم بین ۴۷ تا ۵۸ درصد و الگوریتم بیز ساده از دور پنجم به میزان ۷ درصد به سمت بالا بهبود یافته‌اند. دقت به‌دست‌آمده در روش ترکیبی تک‌واژه‌ای و دوواژه‌ای در الگوریتم‌های ماشین بردار پشتیبان به میزان ۳۰ درصد و کاهش گرادیان تصادفی به میزان ۲۹ درصد و رگرسیون لجستیک به میزان ۲۶ درصد از دور اول و الگوریتم‌های بیز ساده به میزان ۲۰ درصد و جنگل تصادفی به میزان ۳۵ درصد از دور چهارم به سمت بالا بهبود یافته‌اند. دقت به‌دست‌آمده در روش ترکیبی دوواژه‌ای و سه‌واژه‌ای در همه الگوریتم‌های یادگیری ماشین از دور چهارم بین ۲۵ تا ۳۲ درصد به سمت بالا بهبود یافته‌اند. دقت به‌دست‌آمده در روش ترکیبی سه تکنیک در الگوریتم‌های ماشین بردار پشتیبان و کاهش گرادیان تصادفی و رگرسیون لجستیک از دور اول به میزان ۳۰ درصد و الگوریتم‌های بیز ساده به میزان ۲۱ درصد و جنگل تصادفی به میزان ۳۵ درصد از دور چهارم به سمت بالا بهبود یافته‌اند.



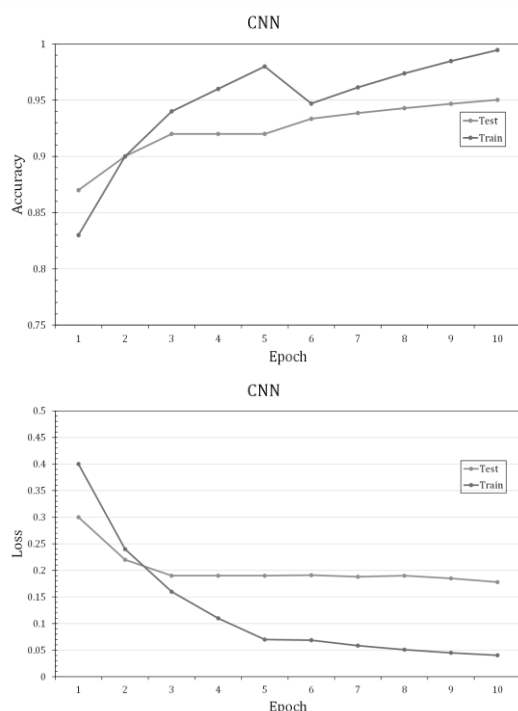
بیشتری بر نتیجه نهایی دارد. مقدار معیار فیشر به‌دست‌آمده در الگوریتم‌های ماشین بردار پشتیبان و کاهش گرادیان تصادفی و رگرسیون لجستیک در همه روش‌ها به‌غیر از روش دوواژه‌ای و سه‌واژه‌ای و ترکیب دوواژه‌ای و سه‌واژه‌ای به میزان ۱۰ درصد بهتر از مقدار حاصل‌شده در الگوریتم‌های دیگر هستند. برای درک بیشتر نمودار آماری الگوریتم‌های مختلف از جمله میانگین، انحراف معیار و چارک‌ها در شکل (۱) آمده است.



(شکل - ۱). نمودار آماری الگوریتم‌های مختلف از جمله میانگین، انحراف معیار و چارک‌ها (دیجی کالا)

(Figure - 4). Accuracy result of execution rounds using different n-gram technics (Digikala)

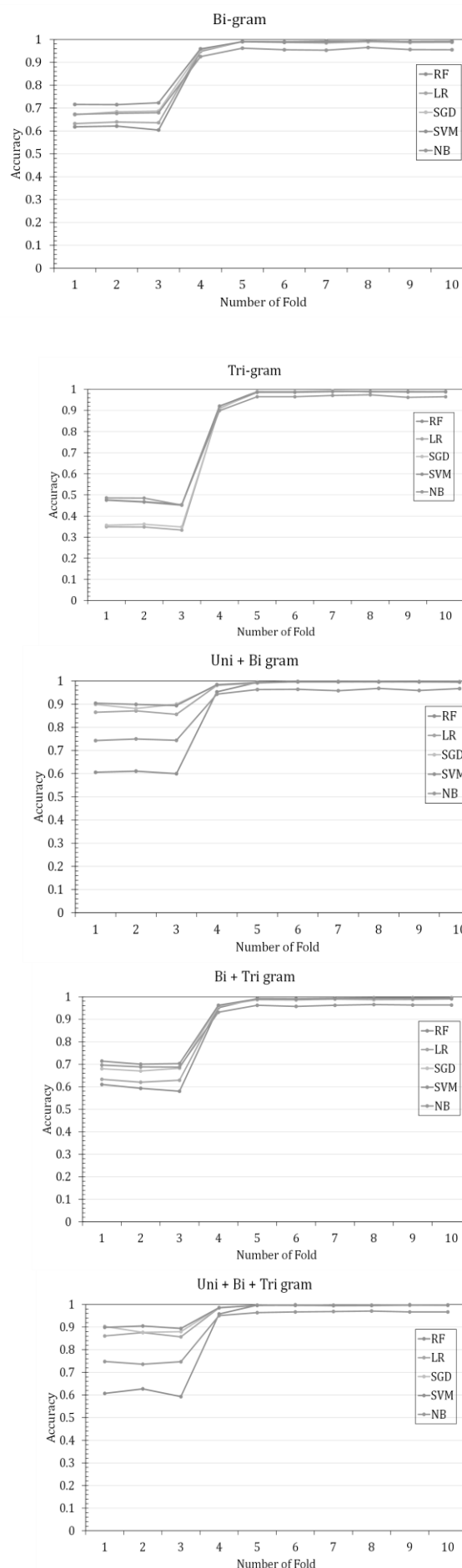
دقت و خطای به دست آمده در روش شبکه عصبی کانولوشن از روش های یادگیری عمیق پس از اعمال روش پیشنهادی در **Error! Reference source not found.** (۵) نشان داده شده است. می توان تحلیل کرد که مقدار دقت به دست آمده بین آموزش و آزمایش تا حدودی از بین رفته است و هرچند در دور سوم به میزان ۲ درصد افزایش دقت در داده های آموزش داریم ولی در داده های آزمایش ثابت شده که حالت بیش برآزش پیدا کرده است و دقت از آموزش کمتر شده است. همچنین درصد خطای روش شبکه عصبی کانولوشن به دست آمده در داده های آموزش در هر دوره بین ۴ تا ۲۵ درصد کمتر شده، ولی برای داده های آزمایش در دوره سوم با ۱۹ درصد خطا ثابت شده و بیشتر از داده های آموزش است. به طور کلی روش CNN بر روی داده های مورد نظر و بعد از اعمال روش پیشنهادی نسبت به الگوریتم های جنگل تصادفی و بیز ساده با ۹۲ درصد دقت بهتر عمل کرده است.



(شکل - ۵) . دقت و درصد خطای به دست آمده با روش شبکه عصبی کانولوشن (دیجی کالا)

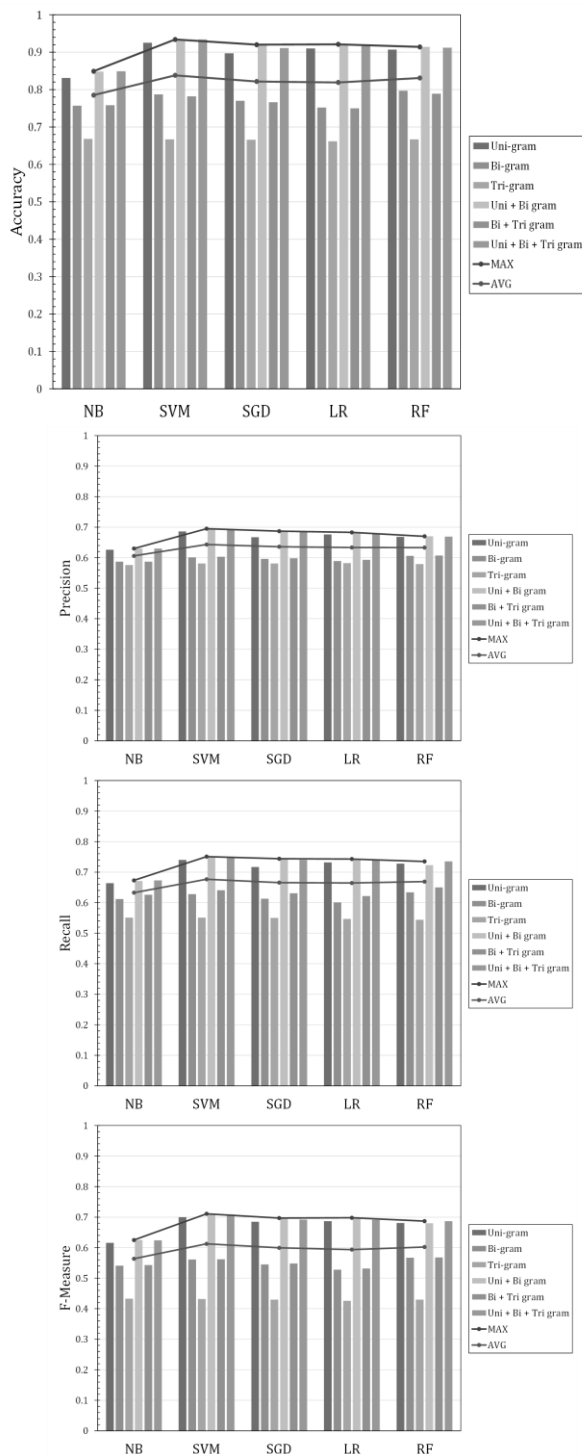
(Figure - 5). Accuracy and loss using Convolution Neural Network (Digikala)

دقت و خطای به دست آمده در روش حافظه کوتاه مدت متوالی از روش های یادگیری عمیق پس از اعمال روش پیشنهادی در **Error! Reference source not found.** (۶) نشان داده شده است. مقدار دقت به دست آمده بین آموزش و آزمایش در همان دور اول هم در آموزش و هم در



(شکل - ۴) نتایج دقت روش بر اساس دوره های مختلف اجرا با استفاده از الگوریتم های یادگیری ماشین (دیجی کالا)

الگوریتم‌های مختلف یادگیری ماشین در Error! Reference source not found. (۷) نشان داده شده است.



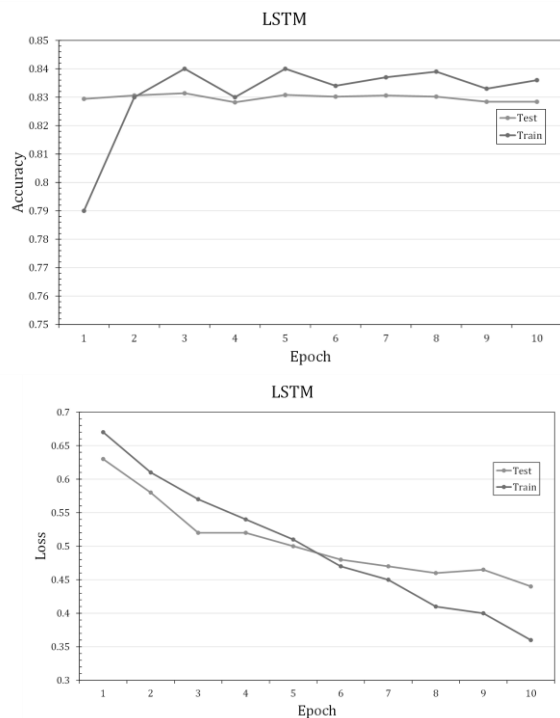
(شکل - ۷). نتایج پارامترهای ارزیابی الگوریتم‌های یادگیری

ماشین بر اساس تکنیک‌های چندواژه‌ای (سنتی پرس)

(Figure - 7). Machine learning algorithms result based on n-gram techniques on Senti-Pers comments

همان‌طور که در Error! Reference source not found. (۲) نشان داده شده است، مقدار دقت به‌دست‌آمده با تکنیک تک‌واژه‌ای در الگوریتم ماشین بردار پشتیبان بین ۱ تا ۹ درصد بهتر از مقدار حاصل‌شده در الگوریتم‌های دیگر است. روش تک‌واژه‌ای به دلیل این‌که به‌صورت

آزمایش با مقدار ۸۳ درصد ثابت شده‌اند و حالت بیش‌برازش پیدا کرده‌اند. هم‌چنین درصد خطای روش حافظه کوتاه‌مدت متوالی در داده‌های آموزش و آزمایش در هر دوره بین ۲ تا ۶ درصد کمتر شده است و در دوره‌های آخر در داده‌های آزمایش درصد خطا پایین‌تر نیز آمده است. به‌طورکلی روش LSTM بر روی داده‌های موردنظر و بعد از اعمال روش پیشنهادی نسبت به الگوریتم‌های یادگیری ماشین و CNN با ۸۳ درصد دقت به دلیل این‌که وابستگی جملات را لحاظ نمی‌کند ضعیف‌تر عمل کرده است.



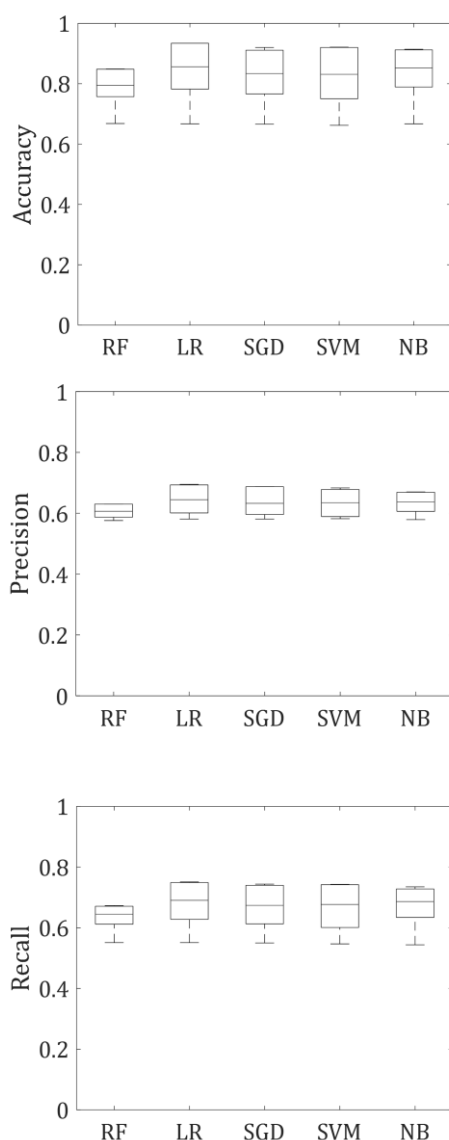
(شکل - ۶). دقت و خطای به‌دست‌آمده در روش حافظه کوتاه‌مدت متوالی (دیجی‌کالا)

(Figure - 6). Accuracy and loss using LSTM Network (Digikala)

۴-۳- تأثیر استفاده از فرهنگ‌نامه و متوازن‌سازی بر داده‌های سنتی پرس

در این قسمت نتایج پارامترهای ارزیابی نظرات تلفن همراه از سایت سنتی پرس با استفاده از الگوریتم‌های یادگیری ماشین توسط تکنیک‌های تک‌واژه‌ای، دوواژه‌ای، سه‌واژه‌ای، ترکیب تک‌واژه‌ای و دوواژه‌ای، ترکیب دوواژه‌ای و سه‌واژه‌ای و ترکیب هر سه تکنیک بررسی شده است. با استفاده از ماتریس درهم‌ریختگی، پارامترهای مختلف ارزیابی مانند دقت، صحت، بازیابی و معیار فیشر به‌دست‌آمده پس از اعمال روش پیشنهادی محاسبه شده است. دسته‌بندی نتایج بر اساس انواع تکنیک‌های چندواژه‌ای و با استفاده از

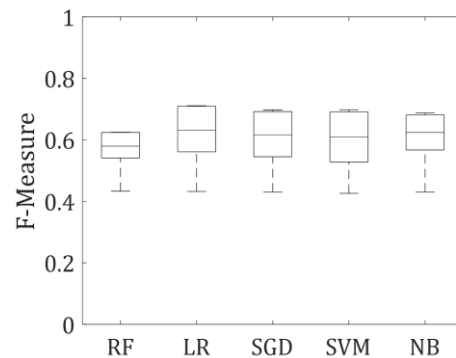
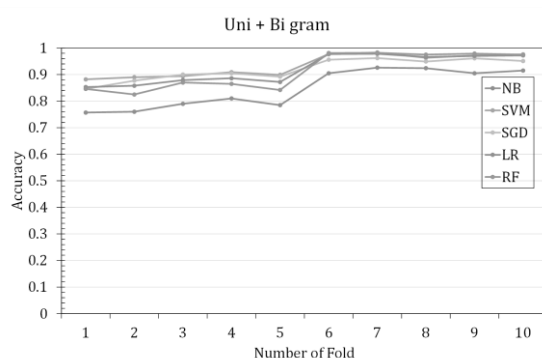
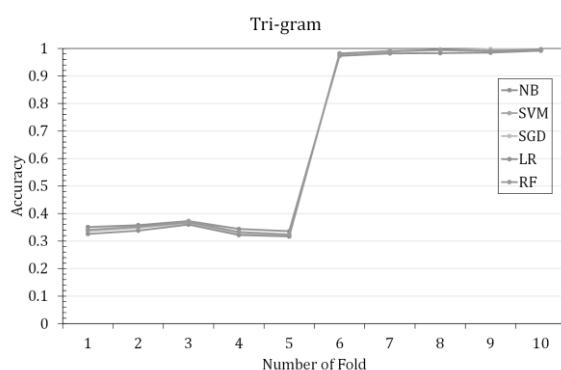
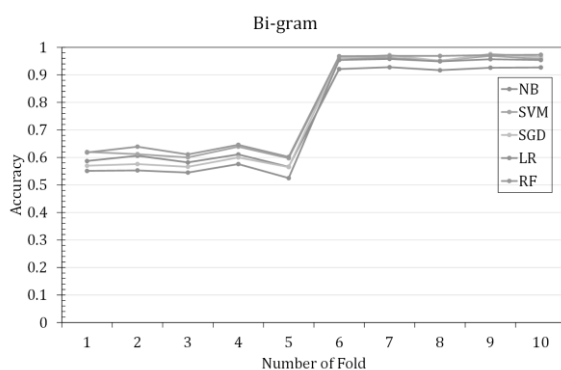
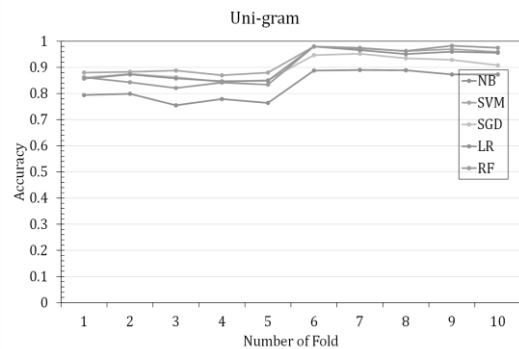
هستند. تأثیر روش سه واژه ای باعث می شود که مقدار دقت نسبت به روش های دیگر در الگوریتم های مختلف نیز نتایج ضعیف تری را به دست آورند. مقدار صحت به دست آمده در الگوریتم ماشین بردار پشتیبان در همه روش ها به غیر از روش سه واژه ای به میزان ۶ درصد بهتر از مقدار حاصل شده در الگوریتم های دیگر در روش های دیگر هستند. همچنین مقدار بازیابی به دست آمده به غیر از روش سه واژه ای در همه الگوریتم های یادگیری ماشین به میزان ۱۹ درصد نتایج بهتری را به دست می آورد. معیار فیشر به دست آمده در الگوریتم های ماشین بردار پشتیبان در همه روش ها به غیر از روش سه واژه ای به میزان ۹ درصد و کاهش گرادیان تصادفی و رگرسیون لجستیک و جنگل تصادفی در همه روش ها به غیر از روش سه واژه ای به میزان ۷ درصد بهتر از مقدار حاصل شده در الگوریتم های دیگر هستند. برای درک بیشتر نمودار آماری الگوریتم های مختلف از جمله میانگین، انحراف معیار و چارک ها در Error! Reference source not found. (۸) آمده است.



کلمه به کلمه پردازش می شود در نتیجه نسبت به روش های دیگر نتایج بهتری را به دست می آورد. مقدار دقت به دست آمده با تکنیک دو واژه ای در الگوریتم جنگل تصادفی بین ۱ تا ۴ درصد بهتر از مقدار حاصل شده در الگوریتم های دیگر است. روش دو واژه ای به دلیل این که به صورت دو کلمه ای پردازش می شود در نتیجه نسبت به روش های دیگر به غیر از تک واژه ای نتایج بهتری را به دست می آورد. همچنین مقدار دقت به دست آمده با تکنیک سه واژه ای در همه الگوریتم های یادگیری ماشین به میزان ۶۶ درصد یکسان هستند. روش سه واژه ای به دلیل این که به صورت سه کلمه ای پردازش می شود و کلمات چند بار تکرار می شوند؛ بنابراین، بر احتمال سند اثر می گذارد در نتیجه نسبت به روش های دیگر دقت دسته بندی را کاهش می دهد. مقدار دقت به دست آمده با تکنیک های تک واژه ای و دو واژه ای در الگوریتم ماشین بردار پشتیبان بین ۱ تا ۹ درصد بهتر از مقدار حاصل شده در الگوریتم های دیگر هستند. روش ترکیبی تک واژه ای و دو واژه ای به دلیل وجود روش تک واژه ای و دو واژه ای در نتیجه نسبت به روش های دیگر نیز نتایج بهتری را به دست می آورند. مقدار دقت به دست آمده با تکنیک های دو واژه ای و سه واژه ای در الگوریتم های ماشین بردار پشتیبان و جنگل تصادفی بین ۲ تا ۳ درصد بهتر از مقدار حاصل شده در الگوریتم های دیگر است. روش ترکیبی دو واژه ای و سه واژه ای به دلیل تأثیر روش سه واژه ای باعث می شود که مقدار دقت نسبت به روش های دیگر نیز نتایج ضعیف تری را به دست آورد. دقت به دست آمده با سه تکنیک در الگوریتم ماشین بردار پشتیبان بین ۱ تا ۹ درصد بهتر از مقدار حاصل شده در الگوریتم های دیگر هستند. روش ترکیبی تک واژه ای و دو واژه ای و سه واژه ای به دلیل تأثیر روش سه واژه ای باعث می شود که مقدار دقت نسبت به روش های دیگر نیز نتایج ضعیف تر و تأثیر روش تک واژه ای و دو واژه ای نتایج بهتری را به دست آورند؛ بنابراین نتایج خوبی را به ارمغان می آورند.

به طور خلاصه دقت به دست آمده پس از اعمال روش پیشنهادی و استفاده از الگوریتم های مختلف یادگیری ماشین مقدار دقت به دست آمده در الگوریتم های ماشین بردار پشتیبان در همه روش های N-Gram به غیر از روش سه واژه ای به میزان ۴ درصد و رگرسیون لجستیک و جنگل تصادفی در همه روش ها به غیر از دو واژه ای و سه واژه ای و روش ترکیبی دو واژه ای و سه واژه ای به میزان ۲ درصد بهتر از مقدار حاصل شده در الگوریتم های دیگر

تصادفی به میزان ۱۳ درصد از دور ششم به سمت بالا بهبود یافته‌اند.



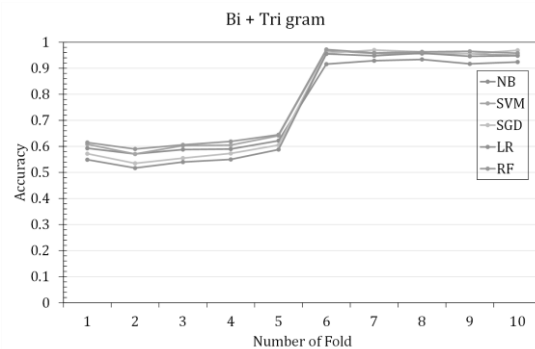
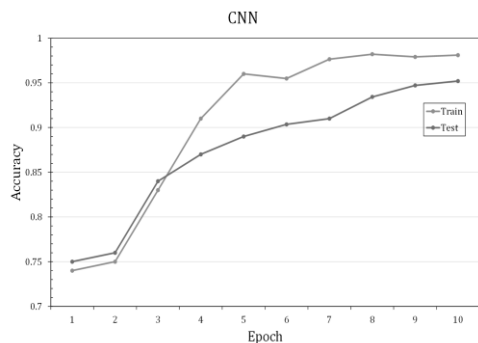
(شکل - ۸). نمودار آماری الگوریتم‌های مختلف از جمله

میانگین، انحراف معیار و چارک‌ها (سنتی پرس)

(Figure - 8). Statistical charts of different algorithms such as Average, SD, and Quarters (Senti-Pers)

دقت به دست آمده پس از اعمال روش پیشنهادی و با استفاده از الگوریتم‌های مختلف یادگیری ماشین با آزمایش در ۱۰ دور متفاوت در **Error! Reference source not found.** (۹) نشان داده شده است. مشاهده می‌شود که مقدار دقت به دست آمده در الگوریتم‌های بیز ساده به میزان ۱۲ درصد و ماشین بردار پشتیبان به میزان ۱۰ درصد و کاهش گرادیان تصادفی به میزان ۹ درصد و رگرسیون لجستیک به میزان ۱۳ درصد و جنگل تصادفی به میزان ۱۵ درصد از دور ششم به سمت بالا بهبود یافته‌اند. دقت به دست آمده در روش دوواژه‌ای پس از اعمال روش پیشنهادی در همه الگوریتم‌های یادگیری ماشین از دور ششم بین ۳۶ تا ۴۳ درصد به سمت بالا بهبود یافته‌اند. همچنین دقت به دست آمده در روش سه‌واژه‌ای در همه الگوریتم‌های یادگیری ماشین از دور ششم بین ۶۴ تا ۶۷ درصد به سمت بالا بهبود یافته‌اند. این مقدار در روش ترکیبی تک‌واژه‌ای و دوواژه‌ای در الگوریتم‌های بیز ساده به میزان ۱۲ درصد و ماشین بردار پشتیبان به میزان ۸ درصد و کاهش گرادیان تصادفی به میزان ۶ درصد و رگرسیون لجستیک به میزان ۱۰ درصد و جنگل تصادفی به میزان ۱۴ درصد از دور ششم به سمت بالا بهبود یافته‌اند. دقت به دست آمده در روش ترکیبی دوواژه‌ای و سه‌واژه‌ای پس از اعمال روش پیشنهادی در همه الگوریتم‌های یادگیری ماشین از دور ششم بین ۲۹ تا ۳۷ درصد به سمت بالا بهبود یافته‌اند. در نهایت دقت به دست آمده در روش ترکیبی سه تکنیک در الگوریتم‌های ماشین بردار پشتیبان به میزان ۹ درصد و کاهش گرادیان تصادفی به میزان ۶ درصد از دور پنجم و بیز ساده به میزان ۱۲ درصد و رگرسیون لجستیک به میزان ۱۱ درصد و جنگل

درصد ثابت شده است و حالت بیش‌برازش پیدا کرده‌اند. هم‌چنین درصد خطای روش حافظه کوتاه‌مدت متوالی در داده‌های آموزش و آزمایش در هر دوره بین ۱ تا ۲ درصد کمتر شده است و در دوره آخر در داده‌های آزمایش درصد خطا با مقدار ۶۲ درصد پایین‌تر از آموزش نیز آمده است. به‌طور کلی روش LSTM بر روی داده‌های موردنظر و بعد از اعمال روش پیشنهادی نسبت به الگوریتم‌های یادگیری ماشین و CNN با ۷۵ درصد دقت به دلیل این‌که وابستگی جملات را لحاظ نمی‌کند ضعیف‌تر عمل کرده است.

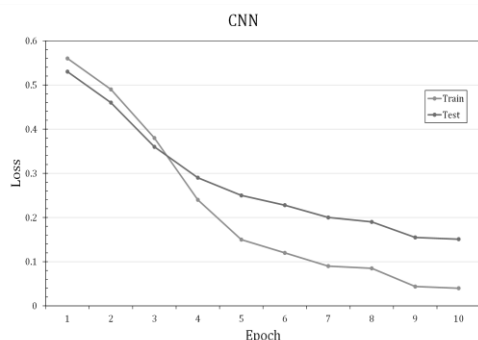


(شکل - ۹). نتایج دقت روش بر اساس دوره‌های مختلف اجرا با

استفاده از الگوریتم‌های یادگیری ماشین (سنتی پرس)
(Figure - 9). Accuracy result of execution rounds using different n-gram techniques (Senti-Pers)

دقت و خطای به‌دست‌آمده در روش شبکه عصبی کانولوشن از روش‌های یادگیری عمیق پس از اعمال روش پیشنهادی در Error! Reference source not found. (۱۰) نشان داده شده است. همان‌طور که نشان داده شده است، می‌توان تحلیل کرد که مقدار دقت به‌دست‌آمده بین آموزش و آزمایش تا حدودی از بین رفته است و هرچند در دور پنجم به میزان ۵ درصد افزایش دقت در داده‌های آموزش داریم ولی در داده‌های آزمایش این مقدار از بین رفته است و بیش‌برازش نمایان شده است. هم‌چنین درصد خطای روش شبکه عصبی کانولوشن در داده‌های آموزش در هر دوره بین ۷ تا ۱۴ درصد کمتر شده است ولی برای داده‌های آزمایش در هر دوره بین ۴ تا ۱۰ درصد خطا کمتر شده است و در دوره آخر با میزان ۲۵ بیشتر از داده‌های آموزش است. به‌طور کلی روش CNN بر روی داده‌های موردنظر و بعد از اعمال روش پیشنهادی نسبت به الگوریتم‌های رگرسیون لجستیک و جنگل تصادفی و بیس ساده با ۸۹ درصد دقت بهتر عمل کرده است.

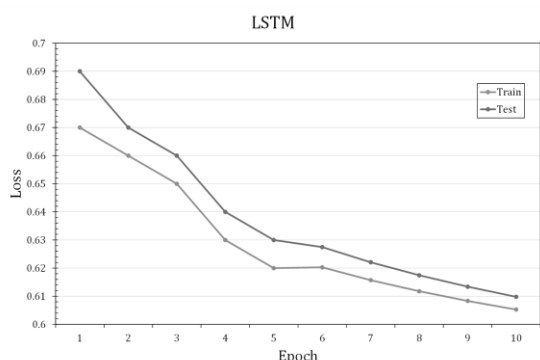
دقت و خطای به‌دست‌آمده در روش حافظه کوتاه‌مدت متوالی از روش‌های یادگیری عمیق پس از اعمال روش پیشنهادی در Error! Reference source not found. (۱۱) نشان داده شده است. مشاهده می‌شود که مقدار دقت به‌دست‌آمده در داده‌های آموزش بین ۲ تا ۱۲ تغییر یافته است ولی در داده‌های آزمایش در دوره سوم با مقدار ۷۵



(شکل - ۱۰). دقت و درصد خطای به‌دست‌آمده با روش شبکه

عصبی کانولوشن (سنتی پرس)

(Figure - 10). Accuracy and loss using Convolution Neural Network (Senti-Pers)



(شکل - ۱۱). دقت و خطای به‌دست‌آمده در روش حافظه

کوتاه‌مدت متوالی (سنتی پرس)

(Figure - 11). Accuracy and loss using LSTM (Senti-Pers)

۴-۴- مقایسه دقت رویکرد پیشنهادی با

کارهای دیگران

نتیجه مقایسه‌ای دقت رویکرد پیشنهادی با کارهای مختلف و استفاده از داده‌های فارسی تلفن همراه از سایت دیجی کالا و داده‌های سنتی پرس از سایت پیکره‌گان در جدول (۲) آمده است. این اطلاعات بر اساس روش‌های مختلف چندواژه‌ای و الگوریتم‌های مختلف یادگیری ماشین دسته‌بندی شده است. اگرچه نویسندگان واقفانند که در مورد مقایسه با کارهای گذشته شاید این مقایسه معنادار نباشد چراکه مجموعه داده‌ها متفاوت است ولی می‌توان با یک نگاه کلی‌نگر تا حدودی نتیجه نهایی را تحلیل نمود. همان‌طور که مشاهده می‌شود، می‌توان تحلیل کرد که مقدار دقت به‌دست آمده با استفاده از روش پیشنهادی روی داده‌های دیجی کالا و سنتی پرس نسبت به روش قبل از اعمال فرهنگ‌نامه و متوازن‌سازی داده‌ها و کارهای دیگران در همه تکنیک‌های چندواژه‌ای و تمام الگوریتم‌های یادگیری ماشین بهتر بوده است. همچنین می‌توان دریافت که دقت به‌دست آمده از روش اعمال پیشنهادی روی داده‌های دیجی کالا در الگوریتم‌های بیز ساده بین ۱۰ تا ۳۴ درصد و ماشین بردار پشتیبان بین ۵ تا ۲۴ درصد و کاهش گرادیان تصادفی بین ۷ تا ۳۸ درصد و رگرسیون لجستیک بین ۵ تا ۳۸ درصد و جنگل تصادفی بین ۴ تا ۲۲ درصد در همه روش‌ها نتایج بهتری را به دست آورده‌اند. همچنین در روش‌های یادگیری عمیق روش CNN با اعمال روش پیشنهادی روی داده‌های دیجی کالا به میزان ۴ درصد بهبود یافته است ولی در روش LSTM هیچ تغییری نیافته است. همچنین دقت به‌دست آمده از اعمال روش پیشنهادی روی داده‌های سنتی پرس در الگوریتم‌های بیز ساده بین ۱۲ تا ۴۶ درصد و ماشین بردار پشتیبان بین ۵ تا ۴۶ درصد و کاهش گرادیان تصادفی بین ۵ تا ۳۵ درصد و رگرسیون لجستیک بین ۶ تا ۴۶ درصد و جنگل تصادفی بین ۴ تا ۴۶ درصد در همه روش‌ها نتایج بهتری را به دست آورده‌اند.

(جدول - ۲) نتیجه مقایسه دقت با ادبیات مختلف و داده‌های

فارسی تلفن همراه از سایت دیجی کالا و داده‌های سنتی پرس
(Table - 2). Accuracy comparison with related works

رویکرد	قبل از اعمال
جدید	روش جدید
Matsumoto	دقیق
Mullen, Collier	دقیق
Beineke	دقیق
Salveti	دقیق
Pang	دقیق

تک‌واژه	۸۱	۷۹.۵	۶۵.۹	-	-	۶۸	۶۸.۴	۸۴.۷	۸۳.۱
دوواژه	۷۷.۳	-	-	-	-	۷۳.۷	۵۴.۴	۸۷	۷۵.۷
سه‌واژه	-	-	-	-	-	۴۷.۷	۳۰.۷	۸۱.۲	۶۶.۸
تک و دوواژه	-	-	-	-	-	۷۹.۴	۷۲.۱	۸۹.۶	۸۴.۸
دو و سه‌واژه	-	-	-	-	-	۷۵.۳	۵۳.۴	۸۷.۷	۷۵.۸
هر سه	-	-	-	-	-	۷۹.۷	۷۱.۵	۸۹.۹	۸۴.۹
تکنیک									
تک‌واژه	۷۲.۹	-	-	۸۶	۸۳.۷	۸۸.۸	۸۷.۳	۹۷	۹۲.۵
دوواژه	۷۷.۱	-	-	-	۸۰.۴	۸۱.۵	۵۴.۴	۹۰.۵	۷۸.۷
سه‌واژه	-	-	-	-	-	۵۸.۱	۳۰.۵	۸۲.۵	۶۶.۷
تک و دوواژه	-	-	-	-	-	۹۰.۷	۸۸.۳	۹۶.۶	۹۳.۴
دو و سه‌واژه	-	-	-	-	-	۸۲.۶	۵۹.۴	۹۰.۲	۷۸.۲
هر سه	-	-	-	-	-	۹۱.۲	۸۸.۵	۹۶.۷	۹۳.۴
تکنیک									
تک‌واژه	-	-	-	-	-	۸۲.۳	۸۳.۵	۹۵.۱	۸۹.۷
دوواژه	-	-	-	-	-	۷۴	۵۷.۳	۸۹.۲	۷۷
سه‌واژه	-	-	-	-	-	۴۱.۲	۳۱.۹	۷۹.۳	۶۶.۶
تک و دوواژه	-	-	-	-	-	۸۴.۷	۸۶.۵	۹۶.۴	۹۲
دو و سه‌واژه	-	-	-	-	-	۷۵.۴	۵۷.۱	۸۹.۱	۷۶.۶
هر سه	-	-	-	-	-	۸۹.۷	۸۶.۵	۹۶.۲	۹۱.۱
تکنیک									
تک‌واژه	-	-	-	-	-	۸۵.۱	۸۳.۳	۹۵.۸	۹۱
دوواژه	-	-	-	-	-	۷۲.۳	۵۵.۲	۸۸	۷۵.۲
سه‌واژه	-	-	-	-	-	۴۰.۸	۳۰	۷۸.۷	۶۶.۲
تک و دوواژه	-	-	-	-	-	۸۹.۶	۸۵.۵	۹۵.۵	۹۲.۱
دو و سه‌واژه	-	-	-	-	-	۷۳.۵	۵۵	۸۷.۸	۷۵
هر سه	-	-	-	-	-	۹۰.۱	۸۶	۹۵.۷	۹۲
تکنیک									
تک‌واژه	-	-	-	-	-	۸۲.۹	۸۶.۹	۸۹.۲	۹۰.۷
دوواژه	-	-	-	-	-	۷۹	۶۰.۱	۸۷.۴	۷۹.۷
سه‌واژه	-	-	-	-	-	۶۰.۱	۳۰.۶	۸۲.۵	۶۶.۷
تک و دوواژه	-	-	-	-	-	۸۳.۵	۸۶	۸۷.۵	۹۱.۴
دو و سه‌واژه	-	-	-	-	-	۸۰	۶۰.۱	۸۷	۷۸.۹
هر سه	-	-	-	-	-	۸۳.۳	۸۶.۷	۸۷.۷	۹۱.۲
تکنیک									
شبکه عصبی	-	-	-	-	-	۸۸	۸۹	۹۲	۸۹
کانولوشن									
حافظه	-	-	-	-	-	۸۳	۷۵	۸۳	۷۵
کوتاه‌مدت									
متوالی									
(LSTM)									

۵- نتیجه‌گیری و کارهای آینده

این پژوهش تلاشی برای دسته‌بندی نظرات فارسی تلفن همراه از سایت دیجی کالا و داده‌های سنتی پرس از سایت پیکره‌گان با استفاده از الگوریتم‌های مختلف یادگیری ماشین تحت نظارت مانند بیز ساده (NB)، کاهش گرادیان تصادفی (SGD)، ماشین بردار پشتیبان (SVM)، رگرسیون لجستیک (LR) و جنگل تصادفی (RF) و از الگوریتم‌های یادگیری عمیق مانند شبکه عصبی کانولوشن (CNN) و حافظه کوتاه‌مدت متوالی (LSTM) است. این الگوریتم‌ها با

- International Joint Conference on Neural Networks (IJCNN). IEEE, 2018.
- [5] Aung, Khin Zezawar, and Nyein Nyein Myo. "Sentiment analysis of students' comment using lexicon-based approach." 2017 IEEE/ACIS 16th international conference on computer and information science (ICIS). IEEE, 2017.
 - [6] Hastie, Trevor, Robert Tibshirani, and Jerome Friedman. "Unsupervised learning." The elements of statistical learning. Springer, New York, NY, 2009. 485-585.
 - [7] Pang, Bo, Lillian Lee, and Shivakumar Vaithyanathan. "Thumbs up? Sentiment classification using machine learning techniques." arXiv preprint cs/0205070 (2002).
 - [8] Salvetti, Franco, Stephen Lewis, and Christoph Reichenbach. "Automatic opinion polarity classification of movie reviews." Colorado research in linguistics 17 (2004).
 - [9] Beineke, Philip, Trevor Hastie, and Shivakumar Vaithyanathan. "The sentimental factor: Improving review classification via human-provided information." Proceedings of the 42nd Annual Meeting of the Association for Computational Linguistics (ACL-04). 2004.
 - [10] Mullen, Tony, and Nigel Collier. "Sentiment analysis using support vector machines with diverse information sources." Proceedings of the 2004 conference on empirical methods in natural language processing. 2004.
 - [11] Dave, Kushal, Steve Lawrence, and David M. Pennock. "Mining the peanut gallery: Opinion extraction and semantic classification of product reviews." Proceedings of the 12th international conference on World Wide Web. 2003.
 - [12] Matsumoto, Shotaro, Hiroya Takamura, and Manabu Okumura. "Sentiment classification using word sub-sequences and dependency sub-trees." Pacific-Asia conference on knowledge discovery and data mining. Springer, Berlin, Heidelberg, 2005.
 - [13] Zhang, Dongwen, et al. "Chinese comments sentiment classification based on word2vec and SVMperf." Expert Systems with Applications 42.4 (2015): 1857-1863.
 - [14] Liu, Shuhua Monica, and Jiun-Hung Chen. "A multi-label Classification based approach for sentiment classification." Expert Systems with Applications 42.3 (2015): 1083-1093.
 - [15] Luo, Banghui, Jianping Zeng, and Jiangjiao Duan. "Emotion space model for classifying opinions in stock message board." Expert Systems with Applications 44 (2016): 138-146.

استفاده از روش های چندواژه ای در مجموعه داده های تلفن همراه به کار می روند. مشاهده می شود که با افزایش مقدار واژه ها، دقت دسته بندی کاهش می یابد، یعنی برای تک-واژه ای و دوواژه ای، نتیجه استفاده از الگوریتم ها بهتر از سه واژه ای است؛ بنابراین الگوریتم های بیز ساده بین ۱۰ تا ۳۴ درصد و ماشین بردار پشتیبان بین ۵ تا ۲۴ درصد و کاهش گرادیان تصادفی بین ۷ تا ۳۸ درصد و رگرسیون لجستیک بین ۵ تا ۳۸ درصد و جنگل تصادفی بین ۴ تا ۲۲ درصد در همه روش های چندواژه ای و روش شبکه عصبی کانولوشن به میزان ۴ درصد روی داده های دیجی-کالا و داده های سنتی پرس در الگوریتم های بیز ساده بین ۱۲ تا ۴۶ درصد و ماشین بردار پشتیبان بین ۵ تا ۴۶ درصد و کاهش گرادیان تصادفی بین ۵ تا ۳۵ درصد و رگرسیون لجستیک بین ۶ تا ۴۶ درصد و جنگل تصادفی بین ۴ تا ۴۶ درصد در همه روش های چندواژه ای عملکرد بهتری نسبت به الگوریتم های دیگر داشتند. از روش های نمونه گیری بیش از حد برای متوازن سازی داده ها و روش امتیازدهی پیشنهادی برای بهبود تعیین قطبیت و از تکنیک های بردار سازی برای تبدیل متن به ماتریس اعداد استفاده شده است و نیز به منظور به دست آوردن مقدار دقت، صحت، بازیابی و معیار فیشر در روش بهبود یافته، از تکنیک های یادگیری ماشین و یادگیری عمیق استفاده می شود. فرایند نظر کاوی در زبان فارسی مشکلاتی از قبیل عامیانه بودن کلمات و رعایت نکردن فاصله بین کلمات و رعایت نکردن قاعده ساختار جمله و استفاده از شکلک های تصویری را دارد که می توان این مشکلات را در کارهای آینده برای بهبود دسته بندی تحلیل احساسات رفع کرد.

6- Refrence

۶- مراجع

- [1] Feldman, Ronen. "Techniques and applications for sentiment analysis." Communications of the ACM 56.4 (2013): 82-89.
- [2] Gautam, Geetika, and Divakar Yadav. "Sentiment analysis of twitter data using machine learning approaches and semantic analysis." 2014 Seventh International Conference on Contemporary Computing (IC3). IEEE, 2014.
- [3] Krizhevsky, Alex, Ilya Sutskever, and Geoffrey E. Hinton. "Imagenet classification with deep convolutional neural networks." Communications of the ACM 60.6 (2017): 84-90.
- [4] Domingues, InLes, et al. "Evaluation of oversampling data balancing techniques in the context of ordinal classification." 2018

- function of several independent variables." *Biometrika* 54.1-2 (1967): 167-179.
- [31] Breiman, Leo. "Random forests." *Machine learning* 45.1 (2001): 5-32.
- [32] Yuan, Yufei, and Michael J. Shaw. "Induction of fuzzy decision trees." *Fuzzy Sets and systems* 69.2 (1995): 125-139.
- [33] LeCun, Yann, et al. "Gradient-based learning applied to document recognition." *Proceedings of the IEEE* 86.11 (1998): 2278-2324.
- [34] Hochreiter, Sepp. "JA1 4 rgen Schmidhuber (1997). "Long Short-Term Memory"." *Neural Computation* 9.8.
- [35] Tripathy, Abinash, Ankit Agrawal, and Santanu Kumar Rath. "Classification of sentiment reviews using n-gram machine learning approach." *Expert Systems with Applications* 57 (2016): 117-126.
- [36] Mouthami, K., K. Nirmala Devi, and V. Murali Bhaskaran. "Sentiment analysis and classification based on textual reviews." 2013 international conference on Information communication and embedded systems (ICICES). IEEE, 2013.
- [37] Zhang, H. "The Optimality of Naive Bayes," In *Proc. Seventeenth Int. Florida Artif. Intell. Res. Soc. Conf. FLAIRS 2004*, vol. 1, no. 2, pp. 1-6. 2004.
- [38] Bishop, Christopher M. *Pattern recognition and machine learning*. springer, 2006.
- [39] Bottou, Léon. "Large-scale machine learning with stochastic gradient descent." *Proceedings of COMPSTAT'2010*. Physica-Verlag HD, 2010. 177-186.
- [40] Menard, Scott. *Applied logistic regression analysis*. Vol. 106. Sage, 2002.
- [41] Breiman, Leo. "Random forests." *Machine learning* 45.1 (2001): 5-32.
- [42] Ketkar, Nikhil, and Eder Santana. *Deep Learning with Python*. Vol. 1. Berkeley, CA: Apress, 2017.
- [43]
- [16] Niu, Teng, et al. "Sentiment analysis on multi-view social data." *International Conference on Multimedia Modeling*. Springer, Cham, 2016.
- [17] Li, Caiqiang, and Junming Ma. "Research on online education teacher evaluation model based on opinion mining." 2012 National Conference on Information Technology and Computer Science. Atlantis Press, 2012.
- [18] Ortigosa, Alvaro, José M. Martín, and Rosa M. Carro. "Sentiment analysis in Facebook and its application to e-learning." *Computers in human behavior* 31 (2014): 527-541.
- [19] Pong-Inwong, Chakrit, and Wararat Songpan Rungworawut. "Teaching senti-lexicon for automated sentiment polarity definition in teaching evaluation." 2014 10th International Conference on Semantics, Knowledge and Grids. IEEE, 2014.
- [20] Wang, Yili, and Hee Yong Youn. "Feature Weighting Based on Inter-Category and Intra-Category Strength for Twitter Sentiment Analysis." *Applied Sciences* 9.1 (2019): 92.
- [21] Mnsefi, Gharizadeh, Sefidsangi, "An overview of opinion mining", The first specialized conference on intelligent computer systems and their applications, Tehran, 2011.
- [22] Noeei, Jalali, Ghaemi, "Opinion mining: An overview of the work done", National Conference on Application of Intelligent Systems (soft computing) in Science and Technology, Quchan, 2013.
- [23] Baccianella, Stefano, Andrea Esuli, and Fabrizio Sebastiani. "Sentiwordnet 3.0: an enhanced lexical resource for sentiment analysis and opinion mining." *Lrec*. Vol. 10. No. 2010. 2010.
- [24] Ngoc, Phan Trong, and Myungsik Yoo. "The lexicon-based sentiment analysis for fan page ranking in Facebook." *The International Conference on Information Networking 2014 (ICOIN2014)*. IEEE, 2014.
- [25] Tang, Huifeng, Songbo Tan, and Xueqi Cheng. "A survey on sentiment detection of reviews." *Expert Systems with Applications* 36.7 (2009): 10760-10773.
- [26] Garreta, Raul, and Guillermo Moncecchi. *Learning scikit-learn: machine learning in python*. Packt Publishing Ltd, 2013.
- [27] McCallum, Andrew, and Kamal Nigam. "A comparison of event models for naive bayes text classification." *AAAI-98 workshop on learning for text categorization*. Vol. 752. No. 1. 1998.
- [28] Hsu, Chih-Wei, Chih-Chung Chang, and Chih-Jen Lin. "A practical guide to support vector classification." (2003): 1396-1400.
- [29] Bottou, Léon. "Stochastic gradient descent tricks." *Neural networks: Tricks of the trade*. Springer, Berlin, Heidelberg, 2012. 421-436.
- [30] Walker, Strother H., and David B. Duncan. "Estimation of the probability of an event as a



مهدیه واحدی پور کارشناس ارشد
مهندسی رایانه گرایش نرم افزار از
دانشگاه صنعتی قم است. زمینه های
پژوهشی وی هوش مصنوعی، داده کاوی
و پردازش متن است.
نشانی رایانامه ایشان عبارت است از:

vahedipoor.m@qut.ac.ir



محبوبه شمسی دکترای خود را در
رشته مهندسی رایانه گرایش نرم افزار
در سال ۱۳۹۰ دریافت کرده است.
ایشان هم اکنون استادیار بخش

مهندسی رایانه دانشکده برق و رایانه در دانشگاه صنعتی قم است. نامبرده تاکنون مقالات علمی بسیاری را در مجلات و کنفرانس های معتبر خارجی و داخلی به چاپ رسانده است. زمینه های پژوهشی وی پردازش تصویر، یادگیری ماشین، یادگیری عمیق و پایگاه داده ها است. نشانی رایانامه ایشان عبارت است از:

shamsi@qut.ac.ir



عبدالرضا رسولی کناری دکترای خود را در رشته مهندسی رایانه گرایش امنیت داده در سال ۱۳۹۰ دریافت کرده است. ایشان هم اکنون استادیار بخش مهندسی رایانه دانشکده برق و رایانه در دانشگاه صنعتی قم است. نامبرده تاکنون مقالات علمی بسیاری را در مجلات و کنفرانس های معتبر خارجی و داخلی به چاپ رسانده است. زمینه های پژوهشی وی امنیت و رمزنگاری، یادگیری ماشین، یادگیری عمیق و کلان داده ها است.

نشانی رایانامه ایشان عبارت است از:

rasouli@qut.ac.ir

