

بهبود فرآیند عقیده‌کاوی به کمک تلفیق روش های گرگ خاکستری و ماشین بردار پشتیبان

چکیده: ظهور وب و رشد مستمر آن، حجم عظیمی از اطلاعات تولید شده توسط کاربران را ایجاد کرده است. در این اطلاعات، به راحتی می‌توان اطلاعات ذهنی ارزشمندی را پیدا کرد، به ویژه در شبکه‌های اجتماعی و پلتفرم‌های تجارت الکترونیک که حاوی اطلاعات مهمی در مورد کاربران هستند. حوزه عقیده‌کاوی در سال‌های گذشته به شدت مورد توجه قرار گرفته است. هر روز مقالات تحقیقاتی جدیدی منتشر می‌شوند که در آنها تکنیک‌های مختلف هوش مصنوعی برای وظایف و برنامه‌های مختلف مرتبط با عقیده‌کاوی اعمال می‌شوند. در این مقاله رویکرد جدیدی مبتنی بر تکنیک ماشین بردار پشتیبان و الگوریتم بهینه‌سازی گرگ های خاکستری برای بهبود فرآیند عقیده‌کاوی پیشنهاد شده است، بگونه‌ای که الگوریتم گرگ های خاکستری برای تعیین ویژگی‌های موثر در فرآیند عقیده‌کاوی به کار رفته است و عملکرد ماشین بردار پشتیبان را بهبود می‌دهد. نتایج پیاده سازی این تحقیق نشان داد که سیستم پیشنهادی توانسته با انتخاب ویژگی‌های موثر به افزایش دقت و پوشاندن خطای تکنیک ماشین بردار پشتیبان کمک کند. سیستم پیشنهادی با استفاده از سه معیار صحت، فراخوانی و دقت مورد ارزیابی قرار گرفت، که صحت برای کلاس اول و دوم به ترتیب برابر با ۰/۶۸ و ۰/۹۲ درصد، فراخوانی برای کلاس اول و دوم به ترتیب برابر با ۰/۹۴ و ۰/۶۳ درصد و دقت ۰/۷۷ درصد بوده است. نتایج حاکی از آن است سیستم پیشنهادی این تحقیق توانسته در هر دو کلاس به نتایج مطلوبی دست پیدا کند.

واژه‌های کلیدی: عقیده‌کاوی، ماشین بردار پشتیبان، الگوریتم بهینه‌سازی گرگ های خاکستری.

Improve opinion mining by combining gray wolf and support vector machine

Abstract: *The advent of the web and its continuous growth has created a tremendous amount of user-generated information. Valuable subjective information is easy to find. Especially on social networks and e-commerce platforms that contain essential information. As a result, the field of belief mining has attracted considerable attention in recent years. New research papers are published every day in which various AI techniques are applied to various tasks and applications related to mining opinion. In this article, a new approach based on the support vector machine technique and the gray wolf optimization algorithm is proposed to improve the opinion mining process. Such that, the gray wolf algorithm is used to determine the practical features in the belief mining process and enhances the performance of the support vector machine. This research showed that the proposed system could increase the accuracy and cover the error of the backup vector technique by selecting practical features. The proposed approach has been evaluated using three criteria; accuracy, recall, and precision. Precision for the first and second classes is 0.68 and 0.92%, respectively, and recall for the first and second classes equals 0.94 and 0.63%, and accuracy was 0.77%. The results indicate that the proposed system of this research achieved favorable results in both categories.*

Keywords: Opinion Mining, Support Vector Machine, Gray Wolf Optimization Algorithm.

۱. مقدمه

بخش مهمی از رفتار جمع‌آوری اطلاعات ما همیشه این بوده است که بفهمیم دیگران چه فکر می‌کنند. با در دسترس بودن و محبوبیت فزاینده منابع سرشار از عقاید فرصت‌ها و چالش‌های جدیدی به وجود می‌آیند که با استفاده از آن‌ها می‌توان از فناوری‌های اطلاعات برای جستجو و درک نظرات دیگران استفاده کرد. فعالیت در حوزه نظرکاوی و تحلیل احساسات، که با پردازش محاسباتی عقیده، احساسات و ذهنیت در متن سروکار دارد، حداقل تا حدی به عنوان پاسخی مستقیم به افزایش علاقه به سیستم‌های جدید رخ داده است که مستقیماً با نظرات به عنوان یک شیء درجه یک سروکار دارند[۱].

نظرکاوی، زیرشاخه‌ای در داده‌کاوی و زبان شناسی محاسباتی است که به تکنیک‌های محاسباتی برای استخراج، طبقه‌بندی، درک و ارزیابی نظرات بیان شده در منابع مختلف خبری آنلاین، نظرات رسانه‌های اجتماعی و سایر محتوای تولید شده توسط کاربر اشاره دارد. تحلیل احساسات اغلب در نظر کاوی برای شناسایی احساسات، عاطفه، ذهنیت و سایر حالات احساسی در متن آنلاین استفاده می‌شود[۲]. عقیده‌ها که بصورت داده‌های متنی جمع‌آوری می‌شوند، دارای اطلاعات ارزشمندی هستند که با توجه زمینه کاربرد، می‌توانند اطلاعات مفیدی را برای افراد و سازمان‌ها تهیه کنند. عقیده‌کاوی بدنبال مشخص نمودن تعیین جهت عقیده‌ها یا به عبارتی دیگر، مثبت و منفی بودن هر عقیده است[۳]. طبقه‌بندی نظرات یا احساسات، مهم‌ترین وظیفه در حوزه ی عقیده‌کاوی است، با توجه به اینکه در سال‌های اخیر، ثبت نظرات مردمی به صورت آنلاین به طور چشمگیری افزایش یافته است. به دست آوردن نتایج عقیده کاوی یکی زمینه‌های اصلی محاسبات احساسات می‌باشد [۴].

الگوریتم‌های متعددی در زمینه عقیده کاوی توسعه یافته که هر یک در بهبود مسئله‌ی عقیده‌کاوی تلاش کرده‌اند. در این مقاله نیز هدف آن است که یکی از جدیدترین روش‌های فراابتکاری ارائه شده را به منظور عقیده کاوی به کار بگیریم. علت انتخاب این

روش طبقه‌بندی، سادگی به کارگیری آن در حل مسائل درک نظرات می‌باشد. رویکرد جدید با ترکیب ماشین بردار پشتیبان و الگوریتم گرگ خاکستری ارائه گردیده که به دنبال عقیده کاوی و تجزیه و تحلیل احساسات موجود در متن‌ها می‌باشد و عقیده ی مثبت و منفی موجود در متن‌ها را مشخص می‌کند. انتخاب ویژگی با الگوریتم گرگ خاکستری سبب کاهش ابعاد و همچنین دقت می‌شود، این الگوریتم با هدف بهبود عملکرد جستجو، افزایش سرعت همگرایی و جلوگیری از گیر افتادن در بهینه محلی ارائه شده است

۲. پیشینه تحقیق

در سال‌های اخیر، مشاهده شده است که عقیده کاوی در سطح وب و رسانه‌های اجتماعی به شکل‌دهی مجدد کسب‌وکارها و تحت تأثیر قرار دادن احساسات و عواطف عمومی کمک کرده است و عمیقاً بر سیستم‌های اجتماعی و سیاسی ما تأثیر گذاشته است. بنابراین جمع‌آوری و مطالعه‌ی نظرات در وب به یک ضرورت تبدیل شده است.

کومار و همکارانش یک استراتژی انتخاب ویژگی مبتنی بر معنا برای انجام وظایف نظرکاوی معرفی کردند. مجموعه ویژگی‌ها، با کمک رویکرد مبتنی بر معنای معرفی شده و به منظور در نظر گرفتن توانایی پیش‌بینی فردی کلمات و به حداقل رساندن ویژگی-های انتخاب شده مورد بررسی قرار گرفتند[۵]. گانش و همکاران از انواع طبقه‌بندی‌کننده‌ها مانند نایویز، جنگل تصادفی، درخت تصمیم، ماشین بردار پشتیبان، K نزدیکترین همسایه، رگرسیون لجستیک برای طبقه‌بندی متن استفاده کردند. در این مقاله طبقه‌بندی‌کننده رگرسیون لجستیک برای طبقه‌بندی متن و ایجاد حس کلمه بالاترین صحت^۱، یادآوری و امتیاز $F1$ را در میان طبقه‌بندی‌کننده‌های دیگر به دست آورد. امتیاز $F1$ میانگین وزنی دقت و یادآوری است. امتیاز $F1$ معمولاً مفیدتر از دقت است، به خصوص اگر توزیع کلاس ناهموار باشد. علاوه بر این طبقه‌بندی کننده نایویز بهترین دقت^۲ را در بین طبقه‌بندی‌کننده‌های مورد

^۱Accuracy

^۲ precision

استفاده به دست آورد [۶]. در مقاله آبدلمینام و همکاران، در بخش اول، از شش روش یادگیری ماشین از جمله درخت‌های تصمیم، رگرسیون لجستیک، نزدیک‌ترین همسایه، جنگل‌های تصادفی، ماشین‌های بردار پشتیبان و نایویز استفاده شد که از روش TF-IDF^۱ (به معنای فراوانی وزنی کلمه کلیدی) برای استخراج ویژگی استفاده شد. بخش دوم شامل آزمایش سه نوع یادگیری عمیق با اعمال جاسازی کلمه به عنوان بردار ورودی بود. نتایج به دست آمده از نظر دقت عملکرد بالاتر و واضحی را برای تکنیک‌های یادگیری عمیق نشان داد [۷]. الزیادی و همکاران دو مدل یادگیری عمیق شامل شبکه‌ی عصبی حافظه‌ی طولانی کوتاه مدت (LSTM^۲) و شبکه‌های عصبی کانولوشن (CNN^۳) پیاده‌سازی کردند. یک مدل ترکیبی از معماری CNN و شبکه عصبی بازگشتی (RNN^۴) ارائه شد که در این مدل ویژگی‌های محلی از طریق CNN به عنوان ورودی RNN برای تحلیل احساسات متون عربی جمع‌آوری شد. طبق نتایج به دست آمده مدل ترکیبی پیشنهادی (CNN-LSTM) عملکرد چشمگیری داشته است [۸].

کیوان پور و همکاران یک روش مبتنی بر واژگان و یادگیری ماشین به نام OMLML^۵ با استفاده از شبکه‌های اجتماعی پیشنهاد کردند. با توجه به روش پیشنهادی، ابتدا قطبیت نظرات نسبت به یک کلمه هدف با استفاده از روشی مبتنی بر واژگان و ویژگی‌های متنی کلمات و جملات تعیین شد. در مرحله بعد، با نگاشت فضای ویژگی در یک بردار ۳ بعدی، نظرات براساس یک روش یادگیری ماشین جدید تجزیه و تحلیل و طبقه‌بندی شدند. نتایج آزمایش‌های کمی و کیفی نشان داد که نگاشت داده‌ها در فضای جدید هزینه آموزش را کاهش می‌دهد و عملکرد روش پیشنهادی به‌ویژه از نظر دقت و زمان اجرا قابل قبول بوده است [۹]. در تحقیق سوزا و همکاران یک الگوریتم جدید برای عقیده کاوی و تجزیه و تحلیل احساسات متن ارسال شده در توییتر ارائه شد که بر اساس خوشه-

بندی بدون نظارت عمل می‌کرد. در این روش از نسخه ترکیبی بهینه‌سازی ازدحام ذرات (PSO) و جستجوی فاخته (CS) استفاده شده بود [۱۰]. در مقاله داو و همکاران یک سیستم توصیه‌گر پیشنهاد شده است که از عقیده کاوی بر اساس تکنیک یادگیری عمیق برای بهبود دقت فرآیند توصیه استفاده کرده است. مدل پیشنهادی شامل دو قسمت می‌باشد. در قسمت ابتدایی، از یک شبکه عصبی پیچیده چند کانالی (MCNN)^۶ برای استخراج بهتر جنبه‌ها و تولید رتبه‌بندی‌های جنبه خاص با محاسبه قطبیت احساسات کاربران در جنبه‌های مختلف استفاده شد. در بخش دوم، رتبه‌بندی‌های جنبه خاص در یک ماشین تجزیه تانسور (TF) برای پیش‌بینی رتبه کلی ادغام شد. نتایج تجربی با استفاده از مجموعه داده‌های مختلف نشان داد که مدل پیشنهادی در مقایسه با روش‌های پایه به پیشرفت‌های قابل توجهی دست یافته است [۱۱]. آخمدوا و همکاران از طبقه‌بندی‌کننده‌های مبتنی بر قانون فازی، شبکه‌های عصبی مصنوعی (ANN) و SVM برای عقیده کاوی استفاده کردند. برای تولید این روش‌ها، یک روش فراابتکاری اصلاح شده برای حل مسائل بهینه‌سازی پارامترهای واقعی یا باینری محدود و غیرمحدود پیشنهاد شد. در این روش از طرح‌های وزن-دهی اصطلاحات مختلف به عنوان تکنیک‌های پیش پردازش داده-ها استفاده شد [۱۲].

کریشنا و همکاران مدل جدیدی برای عقیده کاوی و تحلیل احساسات با بررسی متن‌های ارسال شده در توییتر پیشنهاد کردند. مدل پیشنهادی در این مقاله از تکنیک‌های یادگیری ماشین و رویکرد فازی برای عقیده کاوی و طبقه‌بندی احساسات در متن‌های مورد بررسی استفاده کرده است. [۱۳]. در تحقیق نارایان و همکاران برای تشخیص و بازبینی هرزنامه، یک روش عقیده کاوی پیشنهاد شده است. در این روش از مجموعه‌های مختلفی از ویژگی‌ها، LIWC^۷، تگ‌های POS، ویژگی N-gram^۸ و امتیاز احساسات استفاده شده است. برای طبقه بندی نظرات از شش

^۱ Term Frequency -Inverse Document Frequency

^۲ long short-term memory networks

^۳ convolutional neural network

^۴ recurrent neural network

^۵ opinion mining method based on lexicon and machine learning

^۶ Multi convolutional neural network

^۷ Linguistic Inquiry and Word Count

^۸ An N-Gram is a connected string of N items from a sample of text or speech

تکنیک درخت تصمیم، بیز ساده، SVM، k-نزدیکترین همسایه (KNN)، جنگل تصادفی و رگرسیون لجستیک استفاده شده است [۱۴]. کومار و همکاران بر روی عقیده کاوی وب سایت‌هایی مانند آمازون متمرکز شدند. وب سایت آمازون اجازه می‌دهد تا کاربر به صورت آزاد نظر خود را ارسال کند. سیستم پیشنهادی این تحقیق به طور خودکار، بررسی‌ها را از وب سایت استخراج می‌کند. پس از آن با استفاده از الگوریتم‌هایی مانند طبقه بندی نایو بیز، رگرسیون منطقی نظرات به عنوان مثبت و منفی طبقه بندی می‌کند. در پایان از پارامترهای متریک کیفیت برای اندازه گیری عملکرد هر الگوریتم استفاده شده است [۱۵]. البلتاگی و همکاران یک مدل تجزیه و تحلیل احساسات مبتنی بر تکنیک‌های یادگیری ماشین را که ترکیبی از چندین ویژگی است، ارائه کردند. اکثر آن‌ها با استفاده از احساسات واژگان عربی استخراج شدند. مراحل مختلفی مانند تعیین طول متن، تعداد بخش‌ها در نظر گرفته شد. افزایش دقت به دست آمده در شش مجموعه از هفت مجموعه داده مورد توجه قرار گرفت. این سیستم بر روی مجموعه داده‌های رسانه‌های اجتماعی عربستان، مصر، شام اعمال شد و نتایج قابل قبولی به دست آمد [۱۶].

کریشنا و همکاران مدل جدیدی برای طبقه‌بندی و تحلیل احساسات پیشنهاد کردند. مدل پیشنهادی در این مقاله از تکنیک‌های یادگیری عمیق و رگرسیون بردار پشتیبان برای عقیده‌کاوی و طبقه‌بندی احساسات در متن‌های مورد بررسی استفاده کرده است، پارامترهای مورد استفاده برای ارزیابی، تعداد کلمات بررسی مشتری، دقت، پیچیدگی زمانی و نرخ مثبت کاذب بوده‌اند [۱۷]. شایک و همکاران به تحلیل احساسات و نظرکاوی در مورد داده‌های آموزشی پرداختند. در این تحقیق (۱) نقش تحلیل عاطفی در آموزش را در چهار سطح در نظر گرفته است سطح سند، سطح جمله، سطح نهاد و سطح جنبه، (۲) تکنیک‌های حاشیه نویسی احساسات شامل رویکردهای مبتنی بر واژگان و پیکره محور برای حاشیه نویسی‌های بدون نظارت بررسی گردیده (۳) نقش هوش مصنوعی در تجزیه و تحلیل احساسات با روش‌هایی مانند یادگیری ماشین، یادگیری عمیق، و تبدیل کننده ها مورد

بحث قرار گرفته اند (۴) تاثیر تجزیه و تحلیل احساسات بر رویه‌های آموزشی برای افزایش آموزش، تصمیم‌گیری و ارزیابی ارائه گردیده است [۱۸]. علیان نژادی و همکاران یک روش جدید به منظور بازیابی تصویر مبتنی بر محتوا کردند. برای این منظور، توجه به اهمیت بافت اشیاء در یک تصویر، ویژگی جدیدی تحت عنوان هیستوگرام اختلاف بافت در جهت لبه برابر معرفی شده است. در روش پیشنهادی، ابتدا ویژگی‌هایی شامل ویژگی جدید معرفی شده، از تصاویر آموزشی استخراج شده و سپس تعدادی از این ویژگی‌ها انتخاب شده و با استفاده از این ویژگی‌ها و کلاس هر تصویر و همچنین روش یادگیری ماشین بردار پشتیبان، تصاویر در کلاس‌های مختلف به سیستم آموزش داده شده اند [۱۹]. درختی و همکاران دو مدل شبکه عصبی عمیق شامل شبکه عصبی پیچشی ParsCNN و شبکه عصبی با حافظه بلند کوتاه-مدت دوسویه سلسه‌مراتبی با لایه توجه ParsBiLSTM برای طبقه‌بندی متون فارسی استفاده کردند. کارایی از نظر سه معیار ارزیابی دقت، فراخوانی و مقیاس F-مورد مطالعه قرار گرفت. نتایج آزمایش‌ها نشان می‌دهد که روش ParsCNN میزان دقت ۰/۶۹، فراخوانی ۰/۷۰ همچنین روش ParsBiLSTM میزان دقت ۰/۷۲، فراخوانی ۰/۷۳ و مقیاس F 72/0 دارند که نشان‌دهنده کارایی بالاتر این روش‌ها نسبت به روش‌های طبقه‌بندی متون فارسی مورد مطالعه بوده است [۲۰]. زارع و همکاران یک سیستم تجزیه و تحلیل احساسات بر روی داده‌های تلگرام فارسی پیشنهاد کردند. برای این منظور، چند روش استخراج ویژگی شامل بردار رخداد، فراوانی سند و ماتریس تعبیه کلمات جهت بازنمایی داده‌های متنی به عددی گردید و سپس طبقه‌بندی داده‌ها روش‌های مختلف یادگیری ماشین کلاسیک شامل ماشین بردار پشتیبان، درخت تصمیم، بیز ساده و همچنین روش‌های یادگیری عمیق شامل شبکه عصبی عمیق، شبکه عصبی پیچشی و شبکه‌های حافظه طولانی کوتاه مدت یک‌طرفه و دوطرفه بررسی شد. در نهایت ارزیابی و تحلیل نتایج نشان داد که بهترین کارایی توسط روش استخراج ویژگی ماتریس تعبیه کلمات به همراه شبکه حافظه طولانی کوتاه

مدت دوطرفه با دقت ۶۷/۹۰، صحت ۹۰/۰۱، فراخوان ۸۹/۵۴ و معیار F ، ۷۷/۸۹ درصد به دست آمده است [۲۱].

۳. روش تحقیق

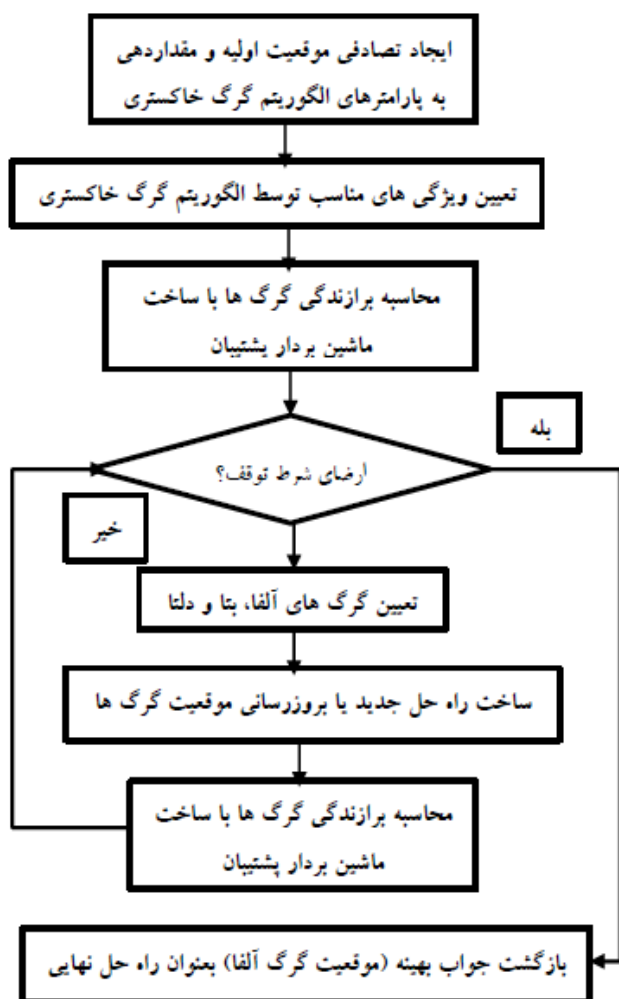
در این تحقیق یک روش ترکیبی مبتنی بر روش‌های ماشین بردار پشتیبان و الگوریتم گرگ‌های خاکستری برای بهبود فرآیند عقیده کاوی ارائه شده است. بخش اصلی سیستم پیشنهادی، ماشین بردار پشتیبان است که الگوهای موجود در داده‌ها را با دریافت داده‌های آموزشی جستجو می‌کند. پس از پایان مرحله‌ی آموزش، سیستم می‌تواند برای عقیده کاوی به کار گرفته شود. برای افزایش کارایی روش پیشنهادی از الگوریتم گرگ‌های خاکستری برای انتخاب مناسب‌ترین ویژگی‌ها استفاده می‌شود به طوری که نتایج به دست آمده مطلوب‌تر خواهد بود. نحوه‌ی ترکیب الگوریتم گرگ‌های خاکستری و ماشین بردار پشتیبان در شکل شماره ۱ آورده شده است.

تعیین ویژگی‌های موثر یکی از مهم‌ترین مسائل در به دست آوردن کارایی مطلوب برای سیستم است. در این الگوریتم، هر فرد نشان‌دهنده تعدادی ویژگی است که در طول تکرارهای مختلف به سمت پاسخ بهینه حرکت می‌کند و در نهایت فردی انتخاب می‌شود که بهترین ویژگی‌ها را برای مسئله انتخاب می‌کند. روشی که برای یادگیری تابع هدف الگوریتم بهینه‌سازی گرگ‌های خاکستری استفاده می‌شود، ماشین بردار پشتیبانی است که از ویژگی‌های مربوط به آن گرگ برای مدل‌سازی استفاده می‌کند.

هدف الگوریتم گرگ‌های خاکستری پیدا کردن مناسب‌ترین ویژگی‌های ممکن می‌باشد، به طوری که ماشین بردار پشتیبان متناظر با آن بهترین عملکرد را داشته باشد.

مطابق با شکل ۱ برای رسیدن به این هدف، در ابتدا موقعیت اولیه‌ی گرگ‌ها به صورت تصادفی ایجاد می‌شوند و به این ترتیب جمعیت اولیه ساخته و پارامترهای آن مقداردهی می‌شود. برای محاسبه‌ی میزان تناسب هر گرگ در سیستم پیشنهادی، ماشین بردار پشتیبان با استفاده از ویژگی‌های انتخاب شده توسط گرگ ساخته می‌شود و دقت ماشین بردار پشتیبان تعیین می‌گردد. در

مرحله‌ی بعد، گرگ‌های آلفا، بتا و دلتا ساخته می‌شوند. در مرحله‌ی بعد موقعیت گرگ‌ها با به روز رسانی موقعیت آن‌ها تغییر می‌کند. پس از آن تناسب هر یک از گرگ‌ها دوباره مشخص خواهد شد. مراحل تعیین گرگ‌های آلفا، بتا و دلتا و به روز رسانی موقعیت گرگ‌ها تا زمانی که شرایط توقف الگوریتم خاتمه یابد، تکرار خواهد شد. در نهایت بهترین راه حل موجود به عنوان نتیجه‌ی نهایی ارائه خواهد شد و ماشین بردار پشتیبان برای آن ساخته شده و با داده‌های آموزش، آموزش می‌بیند. در نهایت دقت به دست آمده از این ماشین بردار پشتیبان به عنوان دقت سیستم پیشنهادی ارائه می‌شود.



شکل ۱: نحوه‌ی ترکیب الگوریتم گرگ‌های خاکستری و ماشین بردار پشتیبان

۱.۳. ماشین بردار پشتیبان

ماشین بردار پشتیبان الگوریتمی برای طبقه‌بندی و دسته‌بندی کردن اشیاء و داده‌ها می‌باشد و مبنای آن برای بهینه‌کردن، کم کردن خطای طبقه‌بندی داده‌های آموزشی می‌باشد [۲۲]. از ویژگی‌های ماشین‌های بردار پشتیبان می‌توان به موارد زیر اشاره کرد:

طراحی دسته‌بندی کننده با ماکسیمم تصمیم، دستیابی به بهینه سراسری تابع هزینه، مشخص نمودن ساختار و توپولوژی دسته‌بندی کننده، پیدا کردن منطقه بهینه توسط مدل‌های غیرخطی.

رویکرد ماشین بردار پشتیبان شکل ویژه‌ای از مدل‌های خطی را مورد استفاده قرار می‌دهد که یک منطقه حاشیه را حداکثر کنند. در واقع به دنبال حداکثر کردن، حداقل فاصله مرز تصمیم‌گیری از طبقه‌ها است. الگوریتم برای یافتن خط جداکننده طبقه‌ها از دو خط موازی استفاده می‌کند که در خلاف جهت هم حرکت می‌کنند تا به یک طبقه برسند، در فاصله بین دو خط یک حاشیه بوجود می‌آید هرچه حاشیه پهن‌تر باشد، یعنی هدف الگوریتم که حداکثر کردن حاشیه بوده است محقق شده است. بیشینه کردن حاشیه ابر صفحه سبب حداکثر شدن جداسازی بین طبقه‌ها می‌گردد. [۲۳]. ماشین بردار پشتیبانی یکی از روش‌های یادگیری بانظارت است که از آن برای طبقه‌بندی خطی و غیرخطی و نیز رگرسیون چندبعدی استفاده می‌کنند. اساس کاری کلاس‌بندی کننده ماشین بردار پشتیبان کلاس‌بندی خطی داده‌ها است و در کلاس‌بندی خطی داده‌ها خطی انتخاب می‌شود که حاشیه اطمینان بیشتری داشته باشد. مساله پیدا کردن خط بهینه برای کلاس‌بندی داده‌ها به وسیله روش‌های برنامه‌ریزی غیرخطی که در حل مسائل محدودیت‌دار شناخته شده هستند صورت می‌گیرد. برای اینکه ماشین بتواند داده‌های با پیچیدگی بالا را نیز دسته‌بندی کند داده‌ها را به وسیله یک تبدیل کرنل به فضای با ابعاد خیلی بالاتر می‌برد.

۲.۳. الگوریتم گرگ خاکستری

الگوریتم گرگ‌های خاکستری با الهام از زندگی اجتماعی و شکار گرگ‌های خاکستری از چهار نوع گرگ برای شبیه‌سازی مراتب سلسله رهبری استفاده می‌کند. گرگ‌های خاکستری یک سلسله

مراتب بسیار سخت اجتماعی دارند که شامل چهار رده آلفا، بتا، امگا و دلتا هستند. در اجتماع گرگ‌های خاکستری موقعیت هر گرگ مشخص است آنها طعمه را محاصره می‌کنند و سپس با دستور رهبر (آلفا) به طعمه حمله می‌کنند. شکار شامل سه مرحله شناسایی، محاصره و حمله به طعمه است. جفت آلفا رهبران گروه هستند که وظیفه تصمیم‌گیری در خصوص مسائل مختلف را به عهده دارند و سایر گرگ‌ها از دستورات وی تبعیت می‌کنند. و به شکلی رفتار نیمه دموکراسی دارند. گرگ‌های آلفا تنها مجاز جفت‌گیری در داخل گروه هستند. آلفا همیشه قوی‌ترین عضو گله نیست، ولی به لحاظ رهبری از سایر گرگ‌ها برتر است و به گرگ‌های ضعیف که توانایی شکار کردن ندارند نیز کمک می‌کند [۲۴].

رده دوم در سلسله مراتب به دسته، بتا است. گرگ بتا، در تصمیم‌گیری‌ها به آلفا کمک کرده و از آن حمایت می‌کند باعث ایجاد نظم و انضباط در گروه می‌شود و بهترین گزینه برای جایگزینی آلفا است. پایین‌ترین سطح در سلسله مراتب گرگ‌ها مربوط به امگا است. آنها ضعیف‌ترین عضو گروه هستند که بیشتر نقش قربانی و پیش‌مرگ را دارند و وظیفه نگهداری از گرگ‌های جوان را دارند. آخرین افرادی در گروه هستند که اجازه خوردن طعمه را دارند. اگر گرگی آلفا، بتا و یا امگا نباشد، زیردست (و یا دلتا در برخی منابع) نامیده می‌شود. گرگ دلتا به لحاظ رتبه در سطح پایین‌تری نسبت به بتا قرار دارند، آنها گرگ‌های قوی هستند ولی قدرت رهبری ندارند و در سطح بالاتری نسبت به امگا قرار دارند. علاوه بر سلسله مراتب اجتماعی، شکار گرگ‌های خاکستری دارای سه مرحله ردیابی، تعقیب و نزدیک شدن به طعمه است. برای مدلسازی سلسله مراتب اجتماعی گرگ‌ها بهترین پاسخ را آلفا و از بین بهترین راه حل‌ها، دومین و سومین را بتا و دلتا در نظر می‌گیریم. بقیه راه حل‌های کاندیدا را امگا در نظر می‌گیریم. بهینه‌سازی توسط آلفا، بتا و دلتا هدایت می‌شود و گروه چهارم از این سه گروه پیروی می‌کند. مدلسازی رفتار محاصره گرگ‌ها از روابط ۱ و ۲ استفاده می‌کند.

$$\vec{D} = |\vec{C} \cdot \vec{X}_p(t) - \vec{X}(t)| \quad (1)$$

$$\vec{X}(t+1) = \vec{X}_p(t) - \vec{A} \cdot \vec{D} \quad (2)$$

حداقل محلی، فاز شناسایی ارائه می‌شود. گرگ‌ها در ابتدا برای شناسایی طعمه از هم دور می‌شوند و بعد از یافتن طعمه و برای حمله به آن به هم نزدیک شده و دور شکار حلقه می‌زنند. برای شبیه‌سازی این واگرایی از بردار A با مقادیر تصادفی بزرگتر از ۱ و یا کوچکتر از ۱- استفاده می‌شود.

جزء تأثیرگذار دیگر بر فرآیند شناسایی مقدار C است. مقدار این بردار عددی تصادفی در بازه $[0, 2]$ است. این مقدار تصادفی تأثیر موقعیت طعمه را در تعیین فاصله، شدت ($C > 1$) یا ضعف ($C < 1$) می‌بخشد [۲۵].

۳.۳. فرآیند ارزیابی در سیستم پیشنهادی

در این بخش، معیارهای مختلفی که برای ارزیابی الگوریتم‌های دسته‌بندی وجود دارند، مورد بحث قرار خواهند گرفت. برای سادگی، تعریف این معیارها برای یک مسئله با دو دسته ارائه می‌شود. در ادامه، ابتدا نگاهی به مفاهیم مثبت واقعی و کاذب، منفی واقعی و کاذب خواهیم داشت. سپس مفهوم ماتریس درهم‌ریختگی^۱ بیان می‌شود. پس از آن تعریف‌هایی از معیارهای مختلف ارزیابی مطرح خواهد شد.

مثبت کاذب^۲ یا FP، پیش‌بینی‌هایی مرتبط با موارد نادرست است که به اشتباه به عنوان درست پیش‌بینی شده‌اند. نرخ مثبت کاذب به کسری از نمونه‌های منفی که به عنوان مثبت تشخیص داده شده‌اند، گفته می‌شود.

		Predicted			
		+	-		
Actual	+	TP Type I error	FN Type II error	Sensitivity (recall) TP/●	False negative rate FN/●
	-	FP Type I error	TN	False positive rate FP/●	Specificity TN/●
		Precision TP/●	False omission rate FN/●	Accuracy (TP + TN)/(● + ●)	
		FDR FP/●	Negative predictive value TN/●	F_1 score $2TP/(2TP + FP + FN)$	

^۱ confusion matrix

^۲ False Positive

در این روابط، t تعداد دفعات تکرار، A و C بردارهای ضریب، $\vec{X}_p$ بردار موقعیت شکار و X بردار موقعیت است. برای محاسبه بردارهای A و C از روابط ۳ و ۴ استفاده می‌شود.

$$\vec{A} = 2\vec{a} \cdot \vec{r} - \vec{a} \quad (3)$$

$$\vec{C} = 2\vec{r}. \quad (4)$$

بردار a از ۲ تا ۰ به صورت خطی در طول دوره تکرار در هر دو فاز اکتشاف و بهره‌برداری کاهش می‌یابد. بردار تصادفی مابین ۰ و ۱ است. با توجه به تصادفی بودن بردارهای r_1 و r_2 ، گرگ‌ها می‌توانند موقعیت خود را در داخل فضایی که طعمه را در برگرفته به صورت تصادفی تغییر دهند. روابط مربوط به آن در معادله ۵ و ۶ آمده است. همین مفهوم را می‌توان به یک فضای جستجوی n بعدی تعمیم داد. در این حالت گرگ‌های خاکستری پیرامون بهترین راه حل به دست آمده در ابعادی بیشتر از ابعاد مکعب حرکت می‌کنند. همانطور که پیش‌تر گفته شد شکار بوسیله آلفا رهبری می‌شود. به منظور مدل نمودن این رفتار سه مورد از بهترین راه‌های به حل دست آمده را ذخیره کرده و دیگر عوامل جستجو طبق رابطه ۷ وادار می‌شوند تا جایگاه خود را با در نظر گرفتن موقعیت بهترین عوامل جستجو و تغییر دهند.

$$\vec{D}_\alpha = |\vec{C}_1 \cdot \vec{X}_\alpha - \vec{X}| \quad (5)$$

$$\vec{D}_\beta = |\vec{C}_2 \cdot \vec{X}_\beta - \vec{X}|$$

$$\vec{D}_\delta = |\vec{C}_3 \cdot \vec{X}_\delta - \vec{X}|$$

$$\vec{X}_1 = \vec{X}_\alpha - \vec{A}_1 \cdot \vec{D}_\alpha$$

$$\vec{X}_2 = \vec{X}_\beta - \vec{A}_2 \cdot \vec{D}_\beta$$

$$\vec{X}_3 = \vec{X}_\delta - \vec{A}_3 \cdot \vec{D}_\delta \quad (6)$$

$$\vec{X}(t+1) = \frac{\vec{X}_1 + \vec{X}_2 + \vec{X}_3}{3} \quad (7)$$

در این الگوریتم پیاده‌سازی فاز بهره‌برداری یا حمله که در صورت متوقف شدن طعمه اتفاق می‌افتد، با کاهش مقدار متغیر a از ۲ به ۱ انجام می‌شود. مقدار A نیز وابسته به a است، از این رو کاهش پیدا می‌کند. با کاهش مقدار A گرگ‌ها مجبور به حمله به سمت طعمه می‌شوند. همچنین برای جلوگیری از گیر افتادن در

۵.۳. معیارهای ارزیابی

از جمله مهم‌ترین معیارهایی که برای ارزیابی الگوریتم‌های دسته‌بندی وجود دارند،

معیارهای

Accuracy

Precision

Recall و F-measure هستند.

معیار دقت یا Accuracy مهم‌ترین معیار برای تعیین کارایی یک الگوریتم دسته‌بندی می‌باشد و دقت کل دسته‌بندی را محاسبه می‌کند [۳]. این معیار با تقسیم تعداد مواردی که به عنوان مثبت شناسایی شده‌اند بر تعداد کل موارد مورد بررسی، بدست می‌آید.

$$Accuracy = \frac{TP + TN}{TP + FP + TN + FN} \quad (12)$$

معیار Precision نشان‌دهنده نسبت کل موارد صحیح طبقه‌بندی شده در طبقه به کل موارد صحیح و غلط در طبقه می‌باشد. این معیار به ما می‌گوید که چند درصد از جواب‌های درست، درست بوده است. به عبارتی صحت دسته‌بندی دسته i را با توجه به کل مواردی نشان. اندیس i در این پارامترها به این مفهوم است که پارامترها باید برای هر دسته i محاسبه شوند [۲۷].

$$Precision_i = \frac{TP_i}{TP_i + FP_i} \quad (13)$$

معیار Recall نسبت موارد واقعی مثبت است که به طور صحیح مثبت پیش‌بینی شده است [۲۸]. در واقع فراخوانی معیاری است که مشخص می‌کند چه اندازه موارد صحیح درست تشخیص داده شده‌اند. حداقل مقدار آن صفر و حداکثر مقدار آن یک می‌باشد. نحوه محاسبه این معیار در ادامه آورده شده است.

$$Recall_i = \frac{TP_i}{TP_i + FN_i} \quad (14)$$

معیار فراخوانی زمانی که ارزش منفی کاذب بالا باشد معیار خوبی است و کارایی دسته‌بندی را نشان می‌دهد و معیار صحت وقتی که ارزش مثبت کاذب بالا باشد معیار مناسبی است و دقت پیش‌بینی را نشان می‌دهد و اینکه مدل چقدر نتیجه را درست پیش

$$False\ positive\ rate = \frac{FP}{FP + TN} \quad (8)$$

منفی کاذب^۱ یا FN پیش‌بینی‌هایی مرتبط با موارد درست است که به اشتباه به عنوان نادرست پیش‌بینی شده‌اند. نرخ منفی کاذب به کسری از نمونه‌های مثبت که به عنوان منفی تشخیص داده شده‌اند، شکل ۲ ماتریس پراکندگی و معیارهای ارزیابی (Lever et al, 2016)

$$False\ negative\ rate = \frac{FN}{TP + FN} \quad (9)$$

مثبت واقعی^۲ یا TP به موارد مثبتی گفته می‌شود که به طور صحیح در کلاس مثبت تشخیص داده شده‌اند. نرخ مثبت واقعی به کسری از نمونه‌های مثبت گفته می‌شود که به درستی به عنوان مثبت طبقه‌بندی شده‌اند.

$$True\ positive\ rate = \frac{TP}{TP + FN} \quad (10)$$

منفی واقعی^۳ یا TN به موارد منفی گفته می‌شود که به طور صحیح در کلاس منفی تشخیص داده شده‌اند. نرخ منفی واقعی به کسری از نمونه‌های منفی گفته می‌شود که به درستی به عنوان منفی طبقه‌بندی شده‌اند [۲۶].

$$True\ negative\ rate = \frac{TN}{TN + FP} \quad (11)$$

۴.۳. ماتریس پراکندگی

ماتریس پراکندگی، تعداد پیش‌بینی‌هایی درست و نادرست را در یک مجموعه داده‌ی شناخته شده نشان می‌دهد. شکل ۲ معرف یک ماتریس پراکندگی است و عناصر آن را شرح می‌دهد.

دایره‌های آبی و خاکستری به ترتیب نشان دهنده‌ی موارد واقعا مثبت (TP + FN) و موارد واقعا منفی (FP + TN) است. مربع‌های آبی و خاکستری نیز، به ترتیب موارد پیش‌بینی شده‌ی مثبت (TP + FP) و پیش‌بینی شده‌ی منفی (FN + TN) را نشان می‌دهد. [۲۶].

^۱ False Negative

^۲ True Positive

^۳ True Negative

بینی کرده است و تا چه اندازه می‌توان به خروجی اعتماد کرد. [۲۹].

معیار F1، یک معیار مناسب برای ارزیابی دقت یک آزمایش (۱۵)

است. این معیار Precision و Recall را با هم در نظر می‌گیرد. معیار F1 در بهترین حالت، یک و در بدترین حالت صفر است. معیار Support یا پشتیبانی نشان دهنده اهمیت مجموعه داده بوده و شاخصی است که تعداد تراکنش‌های یک مجموعه آیت‌م یا اقلام را در پایگاه داده نشان می‌دهد. و به نوعی یک محدودیت تکرار می‌باشد

۴. یافته‌های پژوهش

۴.۱. معرفی مجموعه داده

برای پیاده‌سازی راهکار پیشنهادی، نیاز به مجموعه داده‌ها شامل نظرات افراد می‌باشد. از این رو برای تهیه داده‌های مورد نیاز، از بسترهای آنلاین بهره برده شده است. یکی از این بسترها، وب سایت شناخته شده اسنپ تریپ^۱ می‌باشد که در زیر مجموعه‌های خود، خدمات سفر و نقل و انتقالات را ارائه می‌دهد. به این ترتیب مجموعه نظرات سایت اسنپ تریپ به عنوان مجموعه داده مورد بررسی استفاده شده است همچنین در این پژوهش از زبان برنامه‌نویسی پایتون و کتابخانه‌های هوش مصنوعی مطرح در این زبان بهره‌برده شده است.

۴.۲. آماده سازی داده

به دلیل آنکه متن‌ها اغلب شامل فرمت‌های خاصی که به متن کاوی کمک نمی‌کنند و می‌توانند حذف شوند، لازم است پیش از هرگونه تحلیل بر روی داده‌ها، بسته به نوع مسئله فرمت‌های غیرضروری حذف شوند. فرآیند پیش پردازش در این تحقیق از تکنیک‌هایی چون حذف کلمات ایست واژه‌ها، حذف تک کاراکتر ها، حذف فاصله‌های اضافی، حذف علامت‌های نگارشی، اعمال TFIDF و ایجاد ngram ها استفاده شده است. در فرآیند پیش

پردازش متون انجام شده ابتدا ایست واژه‌ها^۲ حذف می‌شوند. ایست واژه‌ها کلماتی هستند که به طور معمول به شکل گسترده در زبان مورد بررسی مورد استفاده قرار می‌گیرد و با وجود تکرار فراوان در متن از جهت معنایی اهمیت کمی دارند. این کلمات حاوی اطلاعات نیستند مانند ضمایر، حروف اضافه و حروف ربط. علاوه بر این در مراحل بعد تک کاراکتر ها و فاصله‌های اضافی موجود در متن‌ها و علامت‌های نگارشی حذف می‌شوند. پس از آن عملیات TFIDF برای تعیین میزان اهمیت کلمه‌ها نسبت به سند انجام می‌شود، در این مرحله تعداد تکرار کلمه در متن مورد بررسی قرار می‌گیرد هر چه یک کلمه بیشتر در یک متن تکرار شده باشد (TF) ولی در سایر متون کمتر تکرار شود (IDF)، مقدار TF-IDF آن بیشتر می‌شود و این معیار خوبی جهت تشخیص وزن یک کلمه در یک جمله باشد. در واقع نشان می‌دهد که یک کلمه چقدر می‌تواند یکتا و مهم باشد. سپس n-gram ها تعیین می‌شود در مدل N-Gram برای ساخت کوله‌ای از کلمات (BoW)^۳ علاوه بر این که هر کلمه، یک ستون ویژگی (یا بُعد) هست، ترکیب کلمات نیز به ابعاد اضافه می‌شوند. در نهایت با توجه به نامتوازن بودن مجموعه داده‌ها در این تحقیق، داده‌ها متوازن خواهد شد.

۴.۳. اعمال روش ماشین بردار پشتیبان

در این بخش روش ماشین بردار پشتیبان بر روی مجموعه داده‌ی این تحقیق اعمال می‌شود. جدول ۱، نتایج به دست آمده از این بخش را نشان می‌دهد. از اطلاعات موجود در این شکل می‌توان دریافت که با اعمال ماشین بردار پشتیبان بر روی داده‌ها، مقادیر Precision برای کلاس اول و دوم به ترتیب برابر با ۰٫۶۰ و ۰٫۸۶ می‌باشد. مقادیر Recall نیز برای کلاس اول و دوم به ترتیب برابر با ۰٫۹۱ و ۰٫۵۰ می‌باشد. مقدار Accuracy این روش برابر با ۰٫۶۹ درصد می‌باشد.

ماتریس درهم ریختگی به دست آمده حاصل از پیاده‌سازی ماشین بردار پشتیبان در شکل شماره ۳ ارائه گردیده است. با توجه

^۲ Stop Words

^۳ Bag of Word

^۱ snaptrip.com

نشان می‌دهد که روش پیشنهادی نسبت به طبقه‌بندی پایه به نتایج

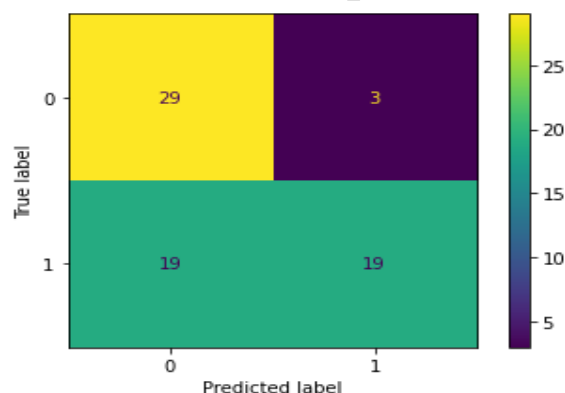
	precision	recall	f1-score	Support
Class 0	0.68	0.94	0.79	32
Class 1	0.92	0.63	0.75	38
Accuracy	0.77			70
macro avg	0.80	0.78	0.77	70
weighted avg	0.81	0.77	0.77	70

بهتری رسیده است.

جدول ۲: نتایج به دست آمده حاصل از پیاده سازی ماشین بردار پشتیبان و الگوریتم گرگ خاکستری

به نتایج حاصله از این ماتریس موارد به اشتباه تشخیص داده شده در دو کلاس به ترتیب برابر با ۱۹ و ۳ است.

جدول ۱: نتایج به دست آمده حاصل از پیاده سازی ماشین بردار پشتیبان



شکل ۳: ماتریس درهم ریختگی به دست آمده حاصل از پیاده سازی ماشین بردار پشتیبان

۴.۴. اعمال روش گرگ خاکستری و ماشین بردار پشتیبان

در این بخش از الگوریتم گرگ‌های خاکستری برای انتخاب ویژگی‌های موثر در مجموعه داده استفاده شده است. در طول اجرای الگوریتم گرگ‌های خاکستری هر راه حل نشان دهنده تعدادی ویژگی است. برای ارزیابی راه حل مورد نظر، ماشین بردار پشتیبان برای ویژگی‌های انتخاب شده پیاده سازی می‌شود. در نهایت بهترین راه حل نشان دهنده ویژگی‌هایی است که مناسب ترین کارایی را دارند. در این بخش، با بهره‌گیری از روش گرگ-های خاکستری مناسب ترین ویژگی‌ها برای استفاده‌ی ماشین بردار پشتیبان تعیین می‌شود.

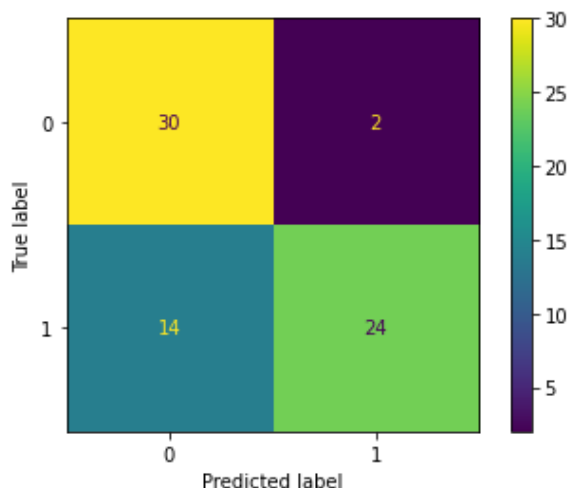
نتایج پیاده سازی روش ماشین بردار پشتیبان با روش گرگ خاکستری در جدول شماره ۲ ارائه گردیده است. از نتایج حاصل می‌توان دریافت که با اعمال سیستم پیشنهادی بر روی داده‌ها، مقادیر Precision برای کلاس اول و دوم برابر با ۰٫۶۸ و ۰٫۹۲ می‌باشد. مقادیر Recall نیز برای کلاس اول و دوم به ترتیب برابر با ۰٫۹۴ و ۰٫۶۳ می‌باشد. مقدار Accuracy این روش برابر با ۰٫۷۷ درصد می‌باشد. مقایسه ی نتایج ثبت شده در این جدول

	precision	recall	f1-score	Support
Class 0	0.60	0.91	0.72	32
Class 1	0.86	0.50	0.63	38
Accuracy	0.69			70
macro avg	0.73	0.70	0.68	70
weighted avg	0.75	0.69	0.68	70

ها پیداست، دقت بدست آمده از روش پیشنهادی، در قیاس با روش دوم (که در آن ویژگی‌های متن که از طریق *CNN* بدست آمده به عنوان ورودی برای شبکه‌های عصبی بازگشتی جهت تحلیل احساسات بکار برده شده) و روش سوم (که بر اساس یک الگوریتم *PSO* خود تطبیقی بهبود یافته با تابع تشخیص و جستجوی فاخته بوده است) و روش چهارم (که از روش‌های بهینه‌سازی محدود باینری و *SVM* استفاده نموده است) دارای دقت بالاتری می‌باشد. در این مقاله، ما یک رویکرد موثر تحلیل عقیده کاوی را با استفاده از معماری مشترک ماشین بردار پشتیبان و گرگ خاکستری (*GOW-SVM*) پیشنهاد کردیم که در مقایسه با برخی روش‌های ترکیبی یادگیری ماشین عملکرد بهتری داشته است. در ادامه پیشنهاد می‌گردد با توجه به عملکرد مناسب مدل اول (یادگیری عمیق و رگرسیون بردار پشتیبان) که نسبت به سایر روش‌های بررسی شده سبب افزایش دقت و کاهش زمان برای استخراج کلمات گردیده است می‌توان از الگوریتم گرگ خاکستری به دلیل کاهش ابعاد و افزایش سرعت همگرایی جهت تعیین ویژگی‌های مناسب پارامترها در *ICSD* استفاده نمود و بدین ترتیب بتوان به عملکرد بالاتری دست یافت.

۵. بحث و نتیجه‌گیری

امروزه اینترنت بشریت را تشویق می‌کند تا نظرات خود را از طریق وبلاگ‌ها، وب سایت‌ها و توییت‌ها منتشر کند. ارزیابی یک سرویس، یک محصول یا حتی یک فرد مشهور و ارسال نظر در مورد هر یک از آنها مفید است. از این رو، طبقه‌بندی نظر کاوی به یک حوزه تحقیقاتی جالب برای تجزیه و تحلیل خودکار این حجم عظیم از نظرات و وبلاگ‌ها تبدیل شده است. برای عقیده کاوی تکنیک‌های مختلفی توسط محققان در گذشته توسعه داده شده است. در این زمینه، رویکردهای مبتنی بر یادگیری ماشین به دلیل مکانیسم یادگیری طبقه‌بندی‌کننده‌ها، امیدوارکننده‌ترین تکنیک در نظر گرفته شده‌اند. در این کار، تکنیک ماشین بردار پشتیبان و رویکرد ترکیبی آن با الگوریتم بهینه‌سازی گرگ‌های خاکستری برای عقیده کاوی مورد بررسی قرار گرفت. نتایج پیاده‌سازی این



شکل ۴: ماتریس درهم ریختگی به دست آمده حاصل از پیاده‌سازی سیستم پیشنهادی

در ادامه نتایج حاصل از روش پیشنهادی به لحاظ دقت با سایر روش‌ها نظیر روش ترکیبی یادگیری عمیق و رگرسیون بردار پشتیبان ^۱ *ICSD* [۱۷]، روش ترکیبی شبکه‌ی عصبی حافظه‌ی طولانی کوتاه‌مدت ^۲ (*LSTM*) و شبکه‌های عصبی کانولوشن ^۳ (*CNN*) [۸]، روش ترکیبی بهینه‌سازی ازدحام ذرات ^۴ (*PSO*) و جستجوی فاخته ^۵ (*CS*) [۱۰]، و روش ترکیبی الگوریتم متاهیوریستیک ^۶ *COBRA* و الگوریتم ^۷ *SVM* [۱۲]، مقایسه گردید که نتایج در جدول شماره ۳ قابل رویت می‌باشد.

جدول ۳: مقایسه نتایج روش پیشنهادی.

روش	دقت (درصد)	
ICSD	78.8	1
LSTM-CNN	70.6	2
IDPSOMUT/CS	63.3	3
COBRA-SVM	62	4
Proposed Method	77	

همانگونه که از جدول مقایسه نتایج روش پیشنهادی با سایر روش

^۱ Lernt Independent Component Support Vector Regressive Deep

^۲ Long Short-term Memory Networks

^۳ Convolutional Neural Network

^۴ Particle Swarm Optimization

^۵ Cuckoo Search

^۶ Co-operation of Biology Related Algorithms

^۷ Support Vector Machine

- pp. 97079-97099, 2021, doi: 10.1109/ACCESS.2021.3094173.
- [8] H. Elzayady. K.M. Badran. G.I. Salama. "Arabic Opinion Mining Using Combined CNN-LSTM Models". *International Journal of Intelligent Systems & Applications*. Vol.12(4), 2020. DOI: 10.5815/ijisa.2020.04.0
- [9] M. Keyvanpour. Z. Karimi Zandian. M. Heidarypanah. "OMLML: a helpful opinion mining method based on lexicon and machine learning in social networks". *Social Network Analysis and Mining*, vol. 10(1). pp. 1-17. 2020. <https://doi.org/10.1007/s13278-019-0622-6>
- [10] E. Souza. D Santos. G. Oliveira. A. Silva. A.L Oliveira. "Swarm optimization clustering methods for opinion mining". *Natural computing*, vol. 19. pp.547-575.2020. <https://doi.org/10.1007/s11047-018-9681-2>
- [11] E. Souza. D. Santos. G .Oliveira. A. Silva, and A.L. Oliveira . "Swarm optimization clustering methods for opinion mining". *Natural computing*. Vol. 19. pp.547-575. 2020. <https://doi.org/10.1007/s11047-018-9681-2>
- [12] A. Da'u. N. Salim. I. Rabiū . A. Osman. "Recommendation system exploiting aspect-based opinion mining with deep learning method". *Information Sciences*. Vol. 512, pp.1279-12. 2020. <https://doi.org/10.1016/j.ins.2019.10.038>
- [13] S. Akhmedova. E. Semenkin. V. Stanovov. "Co-operation of biology related algorithms for solving opinion mining problems by using different term weighting schemes". *Informatics in control, automation and robotics*. p. 73-90. 2018. https://doi.org/10.1007/978-3-319-55011-4_4
- [14] B.V. Krishna. A.K. Pandey. A.S. Kumar. "Feature based opinion mining and sentiment analysis using fuzzy logic". In: *Gurumoorthy S, Rao BNK, Gao XZ (eds) Cognitive science and artificial intelligence*. Singapore, pp 79-89. Springer. 2018. https://doi.org/10.1007/978-981-10-6698-6_8
- [15] R. Narayan. J.K. Rout. S.K. Jena. "Review spam detection using opinion mining". in *Progress in intelligent computing techniques: Theory, practice, and applications*. pp. 273-279. Springer. 2018. DOI:https://doi.org/10.1007/978-981-10-3376-6_30
- سیستم ها نشان داد که تلفیق روش های مذکور می تواند نتایج مناسبی را به دنبال داشته باشد. در سیستم پیشنهادی این تحقیق الگوریتم گرگ های خاکستری برای تعیین ویژگی های موثر در فرآیند عقیده کاوی به کار رفته است و عملکرد ماشین بردار پشتیبان را به نحو مطلوبی بهبود داد به گونه ای که عملکرد سیستم پیشنهاد در هر دو کلاس بهتر از گذشته شد.
- ### منابع
- [1] R.K. Bakshi. N. Kaur. R. Kaur. & G. Kaur. "Opinion mining and sentiment analysis". *3rd International Conference on Computing for Sustainable Global Development (INDIACom)*. 2016. pp452-455.
- [2] H. Chen. D. Zimbra. "AI and opinion mining". *IEEE Intelligent Systems*. 25(3). pp.74-80. 2010. <https://doi.org/10.1109/MIS.2010.75>
- [۳] علی مردانی، س. آقایی، ع. "ارائه روش نظارتی برای نظرکاوی در زبان فارسی با استفاده از لغت نامه و الگوریتم SVM". مدیریت فناوری اطلاعات. جلد ۷، شماره ۲، ص ۳۴۵-۳۶۲، ۱۳۹۴. <https://doi.org/10.22059/jitm.2015.53995>
- [4] M.Y. Raut. M.A. Kulkarni. "Polarity shift in opinion mining". In *2016 IEEE International Conference on Advances in Electronics, Communication and Computer Technology (ICAECCT)*. dec.2016. pp 333-337.
- [5] R.S. Kumar. A.F. Saviour Devaraj. M Rajeswari. E.G. Julie. Y.H. "Robinson . Exploration of sentiment analysis and legitimate artistry for opinion mining". *Multimedia Tools and Applications*. pp.1-16. 2021. <https://doi.org/10.1007/s11042-020-10480-w>
- [6] R. Ganesh. D. Saketh. VDV. S.K. M. Ramesh. "Website Evaluation Using Opinion Mining". In *2021 Third International Conference on Inventive Research in Computing Applications (ICIRCA)*. pp. 1499-1506. IEEE.sep.2021. DOI: [10.1109/ICIRCA51532.2021.9544566](https://doi.org/10.1109/ICIRCA51532.2021.9544566)
- [7] D. S. Abdelminaam, N. Neggaz, I. A. E. Gomaa, F. H. Ismail and A. A. Elsayy, "ArabicDialects: An Efficient Framework for Arabic Dialects Opinion Mining on Twitter Using Optimized Deep Neural Networks," in *IEEE Access*, vol. 9,

