

مدیریت احساسات مشتریان در رسانه های اجتماعی جهت بهبود تبلیغات و افزایش خرید

لیلا خواجهوند^۱، عباس طلوعی اشلقی^{۲*}، مرتضی موسی خانی^۳

^۱ دانشجوی کارشناسی ارشد، گروه مدیریت فناوری اطلاعات، واحد علوم و تحقیقات، دانشگاه آزاد اسلامی، تهران، ایران

^۲ استاد، گروه مدیریت صنعتی، واحد علوم و تحقیقات، دانشگاه آزاد اسلامی، تهران، ایران (عهده دار مکاتبات)

^۳ استاد، گروه مدیریت دولتی، واحد علوم و تحقیقات، دانشگاه آزاد اسلامی، تهران، ایران

تاریخ دریافت: اسفند ۱۴۰۱، اصلاحیه: فروردین ۱۴۰۲، پذیرش: تیر ۱۴۰۲

چکیده:

امروزه شبکه‌های اجتماعی توجه ویژه‌ای را به خود جلب نموده‌اند. در شبکه‌های اجتماعی گوناگون، کاربران دائماً در حال ابراز نظرات عمومی و همچنین خصوصی خود درباره‌ی موضوعات مختلف هستند. توییت‌ر یکی از این شبکه‌های اجتماعی است که در دهه اخیر محبوبیت بسیاری یافته است. این شبکه اجتماعی روشی سریع و موثر برای تحلیل احساسات، دیدگاه‌ها و انتقادات مشتریان برای موفقیت در بازار را به سازمان‌ها ارائه می‌دهد. تحلیل احساسات یا عقیده کاوی فرآیندی است که در آن نظرات، احساسات و نگرش افراد در ارتباط با موضوعی خاص استخراج می‌شود. پژوهش‌های زیادی در رابطه با تحلیل احساسات بر روی نظرات کاربران، مستندات و مقالات انجام شده است. تحلیل بر روی موارد بیان شده تفاوت عمده‌ای با داده‌های توییت‌ر دارد، به این سبب که توییت‌های توییت‌ر محدودیت ۲۸۰ کاراکتری دارند و کاربران را وادار به بیان احساسات خود به صورت فشرده و کوتاه می‌نمایند. بهترین نتایج به دست آمده در طبقه‌بندی احساسات از تکنیک‌های یادگیری ماشین مثل بیز ساده و ماشین بردار پشتیبان حاصل شده است.

در این پژوهش به ارائه روشی برای تحلیل احساسات در شبکه‌های اجتماعی پرداخته می‌شود. در این راستا سعی شده با تمرکز بر مراحل پیش پردازش داده‌ها و انتخاب ویژگی، طبقه‌بندی متن توسط روش بیز را تا حدودی بهبود بخشیم. به عبارتی، با تعریف این مسئله به صورت یک مسئله کلاس‌بندی باینری بر اساس خصیصه‌های پیشنهادی به تحلیل احساسات کاربران پرداخته می‌شود. مسئله کلاس‌بندی با استفاده از جدیدترین دستاوردهای حوزه یادگیری ماشین فرموله و حل شده است. برای ارزیابی روش پیشنهادی در این رساله از سناریو مجموعه دادگان توییت‌ر می‌باشد. روش پیشنهادی با سایر روش‌های طبقه‌بندی مقایسه می‌شود. بهترین عملکرد را از خود نشان داده است.

واژه‌های اصلی: شبکه‌های اجتماعی، تحلیل احساسات، محتوا، یادگیری ماشین، کاربران.

۱- مقدمه

معتقدند که کارخانجات باید تعامل بهتری از طریق رسانه‌های اجتماعی با مشتریان خود داشته باشند [۱۳]. این آمار و ارقام نشان دهنده این موضوع است که چقدر رسانه‌های اجتماعی و ارائه سرویس بر روی این سامانه‌ها برای مشتریان اهمیت دارد و کسب و کارها برای افزایش سهم بازار و پیشرو شدن در رقبا بر روی این پروژه‌ها سرمایه‌گذاری می‌کنند. استفاده از رسانه‌های اجتماعی برای فرآیندهای مرتبط با مشتریان، که به آن مدیریت ارتباط با مشتری اجتماعی گفته می‌شود، بخش بزرگ‌تری از استفاده صرف کسب و کارها از رسانه‌های اجتماعی را شامل می‌شود. در اکثر منابع، مدیریت ارتباط با مشتری به چهار شکل کاملاً متفاوت مورد توجه قرار گرفته است: استراتژیک، عملیاتی، تحلیلی و تعاملی. در نگاه استراتژیک به مدیریت ارتباط با مشتری، CRM یک استراتژی محوری و کلیدی در کسب و کار است و قرار است مشتریان سود ده را جذب کسب و کار کرده و آن‌ها را برای ما حفظ کند.

رسانه‌های اجتماعی با سرعتی تصاعدی در حال گسترش هستند و در حال حاضر آن‌ها فرصتی بی‌نظیر برای کسب و کارها ایجاد کرده‌اند. فیس‌بوک، به‌عنوان بزرگ‌ترین شبکه اجتماعی در حال حاضر بیش از ۹۰۰ میلیون کاربر فعال دارد بیش از ۷۰٪ این کاربر به‌صورت روزانه وارد این سایت شده و در حدود ۹۴۰ میلیارد دقیقه را صرف بازدید از این سایت می‌نمایند. همین وضعیت برای دیگر رسانه‌های اجتماعی نظیر توییتر صادق است. با شیوع بیماری کرونا و همه‌گیر شدن آن، کاربران از صنایع مختلف انتظار دارند سرویس و محصولات خود را در بستر آنلاین به فروش برسانند و آن‌ها بتوانند در هر شرایطی به محصولات مورد نیاز خود دسترسی داشته باشند تحقیق کن ۲۰۲۰ بیش از ۹۸٪ کاربران آمریکایی رسانه‌های اجتماعی عقیده دارند که شرکت‌ها بهتر است خدمات خود را در قالب رسانه‌های اجتماعی ارائه نمایند و بیش از ۸۵٪

*edu.myresearch@gmail.com

در بخش دوم به ارائه کارهای مرتبط پرداخته می‌شود. در بخش سوم کار پیشنهادی مطرح می‌شود و در نهایت در بخش چهارم به ارزیابی روش پیشنهادی پرداخته می‌شود.

۲- پیشینه تحقیق

بیشتر مطالعات مرتبط با تحلیل احساس در گذشته بر اساس الگوریتم‌های یادگیری با نظارت انجام گرفته است که نیاز به تهیه داده برچسب خورده دارند. مدل بیز ساده، ساده‌ترین و پراستفاده‌ترین الگوریتم احتمالاتی برای دسته‌بندی است و بر مبنای قضیه بیز کار می‌کند. این مدل احتمالات پسین رویدادها را محاسبه کرده و برچسبی که بیشترین احتمال پسین را دارد به رویداد نسبت می‌دهد. دسته‌بندی پرکاربرد آنروپی بیشینه است که کار دسته‌بندی را می‌توان با آن انجام داد. این روش بر پایه مدل نمایی و اصل حداکثر آنروپی است. استفاده از این روش تجربه‌های موفق در کار پردازش زبان طبیعی از جمله در تحلیل احساس به ارمغان آورده است. این روش در اکثر (و نه در همه) مواقع نسبت به مدل بیز ساده برتری دارد. ماشین بردار پشتیبان کار دسته‌بندی اسناد بر مبنای موضوعات مشابه بسیار مفید است. روش SVM یک مدل یادگیری بانظارت است که کار آن دسته‌بندی کردن اشیاء در کلاسهای مختلف با استفاده از ویژگی‌های استخراج شده است. این دسته‌بندی با ایجاد ابرصفحه‌ای میان نمونه‌های هر کلاس و حداکثر کردن فاصله نمونه‌ها از این صفحه صورت می‌گیرد. برتری این روش نسبت به دیگر روشهای مطرح یادگیری ماشین آن است که در مورد داده‌های ورودی پیش فرضی ندارد و به جای تکیه بر ارزش‌های احتمالاتی، سعی دارد تا بهینه‌ترین دسته‌بندی را با داده‌های موجود انجام دهد و نتایج به دست آمده از آن در تحلیل احساس برتری محسوس به دیگر روشهای یادگیری ماشین در زبان انگلیسی دارد [۸]. در سالهای اخیر روشهای یادگیری عمیق به خصوص شبکه‌های (RNN) در تحلیل احساس برای زبان انگلیسی چینی و آلمانی در میان زبانهای مختلف، با استفاده از بردارهای مختلف نمایش کلمات کاربرد زیادی داشته است. آنها برای درک و کنترل ترکیب معنایی در کارهای پیچیده‌ای مانند تحلیل احساس مفید هستند. شبکه‌های RNN برای داده‌هایی با قابلیت تبدیل به مقادیر متوالی به کار می‌روند و با استفاده از ایده اشتراک‌گذاری پارامترها برای رسیدن به وزنهای مطلوب، توانایی پردازش توالی‌هایی با طولهای متفاوت را دارند. با وجود این که استفاده از آنها در تحلیل احساس برای زبان انگلیسی با نتایجی بهتر از روشهای یادگیری بانظارت همراه بوده است. با رشد ساختار شبکه‌های RNN، ابعاد ماتریس‌ها در مرحله بازپخش به صورت توانی رشد می‌کنند و در عمل استفاده از آنها غیر ممکن می‌شود [۱۱]. شبکه‌های پیچشی که کولوبرت و دیگران در ابتدا برای کاربرد در بینایی رایانه‌ای ارائه کرده‌اند، اخیراً در بسیاری از کارهای پردازش زبان طبیعی مانند تجزیه نحوی، تجزیه سطحی، برچسب‌زنی نقش معنایی مورد استفاده قرار گرفته است. استفاده از

در نگاه عملیاتی به مدیریت ارتباط با مشتری، هر فرایندی که به‌نوعی با مشتری مرتبط است. با استفاده از سامانه‌های نرم‌افزاری، به ابزارهای خودکار تجهیز می‌شود. بازاریابی، فروش و خدمات مشتریان، ازجمله این فرایندها هستند. در نگاه تحلیلی به مدیریت ارتباط با مشتری، CRM ابزاری برای تحلیل هوشمندانه داده‌ها و اطلاعات مربوط به مشتری باهدف‌های استراتژیک و یا عملیاتی است. در نگاه تعاملی به مدیریت ارتباط با مشتری، CRM از فناوری برای مدیریت مرزهای سازمان، چه در رابطه با مشتری و چه شرکای تجاری استفاده می‌شود و هر نوع داده و اطلاعاتی که از مرزهای سازمان عبور می‌کند توسط این سیستم مدیریت می‌شود. هدف این سیستم ایجاد ارزش برای مشتریان و همکاران سازمان است. با هر دیدگاهی که به CRM توجه شود رسانه‌های اجتماعی می‌تواند جایگاه مهمی داشته باشد. یکی از زیر بخش‌های مهم برای تحقق SCRM، تحلیل احساسات و علایق مشتریان می‌باشد. با توجه به وسعت و قابلیت‌های تعامل کاربران، این بخش نیازمند تحول بنیادین با توجه به نیازهای روزافزون فضای کسب و کار و روش‌های بازاریابی می‌باشد [۱].

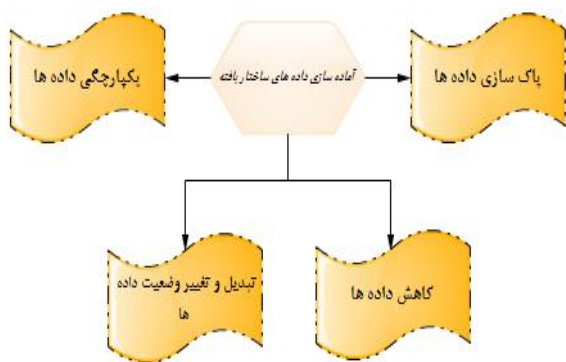
تحلیل احساسات و علایق مشتریان قلب هر سامانه مدیریت ارتباط با مشتری محسوب می‌شود. مسئله مطرح در این پروژه، ارائه روشی برای تحلیل احساسات مشتریان در سامانه‌های مدیریت ارتباط با مشتری اجتماعی است. این روش قادر خواهد بود با بررسی پست‌ها و فعالیت‌های کاربران در رسانه‌های اجتماعی به شناسایی احساسات، اقدام به گروه‌بندی مشتریان بر حسب علایق و احساسات نماید. همین‌طور این روش قادر است گزارش‌گیری کارآمدی از تغییرات نظرات و احساسات مشتریان در بازه‌های زمانی مختلف را ارائه نماید. تحلیل احساس یکی از فعالترین حوزه‌های پژوهشی در پردازش زبان طبیعی است که به دلیل اهمیت آن در کسب و کار و جامعه به خارج از علم کامپیوتر مانند مدیریت و علوم اجتماعی نیز گسترش یافته است. اهمیت درحال رشد تحلیل احساسات با رشد رسانه‌های اجتماعی مانند بررسی‌ها، انجمن‌ها، بحث، وبلاگ‌ها، میکروبلگ‌ها، توییت‌ها و شبکه‌های اجتماعی همخوانی دارد. سیستم‌های سنجش در هر حوزه تجاری و اجتماعی کاربرد دارند؛ زیرا نظرات در همه احساسات تقریباً فعالیت‌های انسانی متمرکز هستند و تأثیرات کلیدی رفتارهای ما به شمار می‌روند. کسب و کارها با تحلیل احساسات کاربران خود می‌توانند سرویس یا محصول خود را ارتقا دهند یا محصولات و سرویس‌های جدید را به بازار ارائه دهند. از طرفی تحلیل احساسات کاربران می‌تواند به تصمیمات مدیران سازمان جهت ارائه تبلیغات موثرتر، بهبود پرموشن‌ها و جذب کردن کاربران کمک کند. همچنین در صورتی که مشتریان نسبت به محصولات یا سرویس‌های فعلی احساس ناخوشایندی داشته باشند می‌توان با دخیل کردن نظرات آن‌ها این حس ناخوشایند را به احساس مثبت و افزایش رضایتمندی آن‌ها تبدیل کرد.

فزاینده‌ای در حال افزایش است. ذات غیرساخت یافتی این متون، اعمال همان روشهایی را که ما در مورد دیتابیس‌ها بکار می‌بریم، غیر ممکن می‌سازد. کاربردهای مهمی را که از پردازش متون مورد انتظار است، بررسی می‌کنیم. به اینگونه پردازش‌ها که روی متون اعمال می‌شود، متن‌کاوی می‌گوییم [۵]. وظیفه‌ی اصلی عقیده‌کاوی طبقه‌بندی قطبیت است. طبقه‌بندی قطبیت وقتی اتفاق می‌افتد که یک تکه متن که یک عقیده در مورد یک موضوع را بیان کند به یکی از دو احساس متضاد تقسیم شود. نظراتی مثل «موافق» در مقابل «مخالف»، «دوست داشتن» در مقابل «دوست نداشتن» مثال‌هایی از طبقه‌بندی عقاید هستند. طبقه‌بندی قطبیت بیانات موافق و مخالف را تشخیص می‌دهد و به تولید ارزیابی‌های معتمدتر کمک می‌کند [۲]. بسیاری از شرکت‌ها از عقیده کاوی و تحلیل احساسات به عنوان جزئی از تحقیقاتشان استفاده می‌کنند. مثلاً شرکت‌ها از عقیده کاوی برای ساخت و نگهداری نظرات استفاده می‌کنند. سیستم‌های آن‌ها به طور مداوم اطلاعات را از وب مثل نظرات محصولات، دریافت برند و مسائل سیاسی جمع‌آوری می‌کند. دیگر سیستم‌ها نیز ممکن است از عقیده کاوی و تحلیل احساسات به عنوان یک فناوری زیر مؤلفه برای بهبود مدیریت روابط مشتری و سیستم توصیه‌گر از طریق بازخوردهای مثبت و مشتریان استفاده کنند. به طور مشابه، عقیده کاوی و تحلیل احساسات ممکن است شعله‌ها (زبان خصومت‌آمیز و گرمای اضافی) را در روابط اجتماعی شناسایی و حذف کنند [۱۲]. طبقه‌بندی بیان احساسات بر اساس معنای آن‌ها و دانش قبلی تمایلات معنایی نامیده می‌شود. با وجود اینکه تحلیل نحوی نقشی کلیدی در طبقه‌بندی اسناد بازی می‌کند اما این برای استخراج مفاهیم از متن فقط از طریق نحو کافی نیست. معیارهای تئوری اطلاعات و دانش معنایی یک سلسله مراتب را با استفاده از WordNet ترکیب کردند تا مفاهیم را به طور اتوماتیک از متن استخراج کنند [۱۵] دنکه نقش مدل بر پایه قوانین و یادگیری ماشین در یک دامنه چندگانه، بر روی سناریو طبقه‌بندی تست کرده است نتایج آن‌ها نشان می‌دهد که رویکرد مبتنی بر واژگان، که با استفاده از SentiWordNet ساخته شده است، دقت آن در مقایسه با روش‌های یادگیری ماشین محدودتر است [۱۰].

۳- روش پیشنهادی

در این قسمت به ارائه روش پیشنهادی برای تحلیل احساسات مشتریان در رسانه اجتماعی خواهیم پرداخت. این روش شامل پنج مرحله به صورت شکل می‌باشد. این مراحل شامل جمع‌آوری داده، پیش پردازش داده، آماده‌سازی داده، برچسب‌گذاری کلمات، شناسایی خصیصه‌های مرتبط با دامنه و خوشه‌بندی کلمات مرتبط با احساسات می‌باشد. این مراحل به صورت تفصیلی در ادامه بخش تشریح خواهد شد.

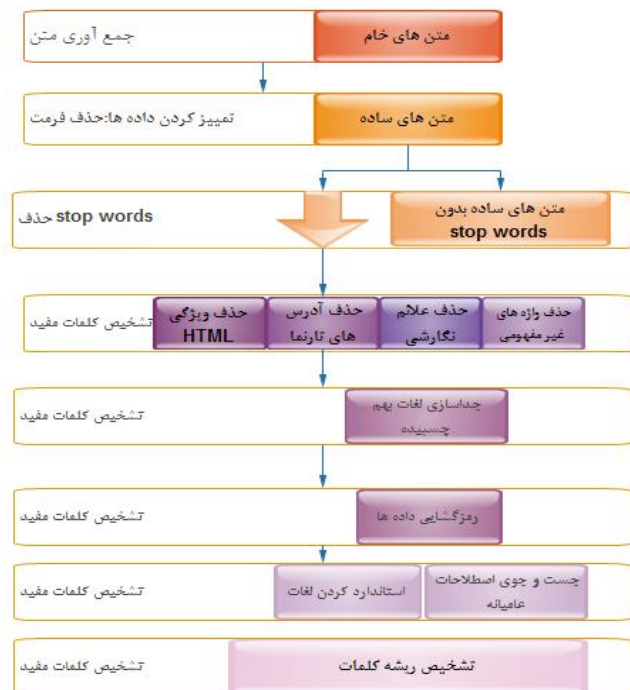
شبکه‌های پیچشی قطعه‌بندی در تحلیل احساس نیز برای زبان‌ها با منابع فراوان مورد استفاده قرار گرفته و باعث بهبود قابل توجه دقت و کاهش زمان مرحله آموزش نسبت به دیگر روشهای یادگیری عمیق شده است. پژوهش‌های حوزه تحلیل احساس در زبان فارسی معمولاً یا با استفاده از روشهای مبتنی بر قاعده هستند یا مبتنی بر پیکره. برای بهبود نتایج معمولاً از پیش‌پردازش نظرات و ویژگی‌های لغتنامه استفاده شده است. بصیری و همکاران [۶] یک چارچوب مبتنی بر لغتنامه ارائه کردند که به صورت بدون نظرات با استفاده از قواعد از پیش تعیین شده و لغتنامه تعریف شده جهت‌گیری متون محاوره را تشخیص می‌دهد. استفاده از SVM برای تحلیل احساس در زبان فارسی بر روی داده مربوط به نقد فیلم، منجر به نتایج بهتری نسبت به روشهای دیگر یادگیری ماشین شده است. بازدهی این روشها وابسته به کیفیت برچسبدهی در پیکره‌ها و شیوه‌گزینش ویژگی‌ها پیش از شروع کار دسته‌بندی است. روشنفکر و همکاران [۷] برای اولین بار از شبکه‌های عصبی LSTM برای تشخیص احساس متون فارسی استفاده کردند و توانستند نسبت به روشهای یادگیری سنتی نتایج بهتری داشته باشند، اما این نوع شبکه‌ها برای آموزش نیاز به داده‌های خیلی زیادی دارند. همچنین آنها در کار خود فقط دو سطح از احساس را در نظر گرفتند و از جاسازی ساده کلمات استفاده کردند. اینترنت از قابلیت‌ها و امکانات زیادی برای ایفای کارکرد حوزه عمومی برخوردار است. شبکه اینترنت امکاناتی در اختیار مردم جوامع می‌گذارد تا در فضایی مناسب به گفتگوی آزاد و برابر با هم بپردازند و در نتیجه فرآیندهای گفتگو و مباحثه، به نقطه‌نظرهای مشترکی درباره مسائل سیاسی و اجتماعی دست یابند و به افکار عمومی شکل دهند. (میناوند، محمدقلی) مجموعه داده‌های نشات گرفته از تحقیق بر روی شبکه‌های اجتماعی در زمینه‌های بسیاری مانند جامعه شناسی و روانشناسی بالارزش هستند. اما حمایت از دیدگاه فنی به اندازه کافی دور است، و به روشهای خاص فوری نیاز دارند. (اکسیو وانگ و همکار) تحقیق روانشناسی برای تشخیص کاربران افسرده شبکه‌های اجتماعی [۱۶] را با استفاده از داده‌کاوی انجام داده است. اخیراً ایده استفاده از تجزیه و تحلیل احساسات کاربران شبکه‌های اجتماعی برای بهبود عملکرد برنامه‌های کاربردی در وب سایت‌های خرید آنلاین توجه پژوهشگران را به خود جلب کرده است. وب سایت‌های خرید آنلاین بطور گسترده‌ای بررسی در مورد یک محصول ارائه می‌کنند و مشتریان می‌توانند استفاده کنند [۴]. در اینجا تجزیه و تحلیل احساسات مشتریان در باره هر محصول انجام می‌شود. در بررسی از سایت‌های شبکه اجتماعی قوانین برای احساسات مثبت یا منفی بسته به نمره کلی آن، با کمک SentiWordNet محاسبه می‌شود. تجزیه و تحلیل تمایلات شامل تشخیص ذهنیت و احساسات موجود در نظرات است. نظرات عبارات توصیف عواطف و احساسات مردم در مورد یک موضوع، نهاد و یا رویداد است با استفاده از تکنیک‌های زبان طبیعی انجام توصیف می‌کنند. [۱۴] تقاضا برای اطلاعات فرابری شده از منابع متنی به طور



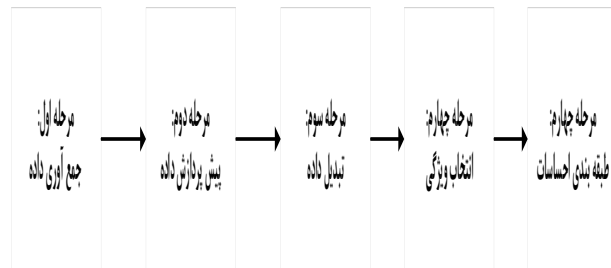
شکل (۲): آماده سازی داده های ساختار یافته

–آماده سازی داده های نیمه ساختار یافته و غیر ساختار یافته

داده های نیمه ساختار یافته شکلی از داده های ساختار یافته ای هستند که از ساختار رسمی از جداول و مدل های داده ای وابسته به پایگاه داده های رابطه ای مطابقت نمی کنند، اما با این وجود شامل برچسب ها یا علامت ها و شاخص هایی هستند که عناصر معنایی را از یکدیگر جدا می کنند و سلسله مراتبی از رکوردها و فیلدها را بین داده ها ایجاد می کنند. داده های غیر ساختار یافته به داده هایی گفته می شود که از هیچ مدل از قبل تعریف شده ای تبعیت نمی کنند مثالی از داده های ساختار یافته متن های سنگین می باشد.



شکل (۳): آماده سازی داده های نیمه ساختار یافته و غیر ساختار یافته



شکل (۱): روش پیشنهادی

مرحله اول، ورودی نظرات کاربران

کاربران با ثبت نام بر روی شبکه های اجتماعی، فعالیت خود را بر بستر این شبکه ها آغاز می کنند. این فعالیت ها شامل تولید محتوا توسط کاربران، پیگیری علاقمندی هایشان، عضو شدن در گروه های متنوع و غیره می باشد. تمامی فعالیت های انجام شده توسط این کاربران در پایگاه داده ذخیره می شود، بنابراین انواع مختلفی از داده های متنی یا غیر متنی، داده هایی با ساختارهای متنوع، داده های غیر دقیق و نادرست وجود دارد. در مراحل بعدی سعی در تمیز کردن داده ها می شود. در مرحله نخست به جمع آوری داده پرداخته می شود. داده هایی می توانند از رسانه های مختلف جمع آوری شود. در این پایان نامه داده تویتر انتخاب شده است. برای جمع آوری داده های تویتر چندین ابزار کاربردی متفاوت، وجود دارد. (۱) برنامه کاربردی جستجوی تویتر، (۲) برنامه کاربردی جریان تویتر، (۳) ابزار آتش نشانی تویتر.

مرحله دوم، پیش پردازش داده ها: جداسازی جملات

مشتریان با ثبت نام بر روی رسانه های اجتماعی، فعالیت های خود را بر بستر این شبکه ها آغاز می کنند. این فعالیت ها شامل تولید محتوا توسط کاربران، پیگیری علاقمندی هایشان، عضو شدن در گروه های متنوع، ارائه نظرات مثبت و منفی خود نسبت به یک محصول یا سرویس خاص می باشد. تمامی فعالیت های انجام شده توسط این مشتریان در پایگاه داده ذخیره می شود، بنابراین انواع مختلفی از داده های متنی یا غیر متنی، داده هایی با ساختارهای متنوع، داده های غیر دقیق و نادرست وجود دارد. در این مرحله برای از بین بردن عدم این ناسازگاری ها، کمبودهای یکپارچگی و بهبود تحلیل نتایج به آماده سازی داده ها پرداخته می شود، برای این کار داده ها، به سه دسته ساختار یافته، نیمه ساختار یافته و غیر ساختار یافته تقسیم می شود.

–آماده سازی داده های ساختار یافته

داده های ساختار یافته از درجه بالای سازمان دهی برخوردار هستند. برای آماده سازی داده های ساختار یافته چهار گام اصلی که در شکل ۲ مشاهده می شود، طی خواهد شد.

مرحله سوم: تبدیل داده

در این بخش داده ها را به فرمتی مناسب تبدیل کرده و آن ها برای مرحله بعدی آماده می شوند.

گاهی داده های خامی که برای تحلیل داریم مناسب گروهی از آزمونهای آماری نیستند و برای این که بتوانیم از این دسته آزمون های آماری استفاده کنیم و همچنین دقت تحلیل را بالا ببریم باید در داده های خام تغییراتی ایجاد کنیم. یکی از این تغییرات، تبدیل داده ها نام دارد. تبدیل داده ها، روش هایی ریاضی است که برای تعدیل متغیرهایی به کار می رود که از مفروضه های آماری نرمال بودن، خطی بودن و یکسانی پراکندگی پیروی نمی کنند یا الگوهایی با داده های پرت غیرمعمول دارند

در مجموع زمانی که پیش شرط های آزمون های چندمتغیره برقرار نباشد، باید داده های به دست آمده را تبدیل کنیم تا امکان استفاده از برخی آزمون های مدنظر (عموما پارامتریک) فراهم شود.

در ابتدا باید میزان تخطی و تفاوت داده ها از پیش فرض های ذکر شده را تعیین کرد و در صورتی که پیش فرض ها یا پیش شرط های آماری به دست آمده دارای تفاوت قابل اعتنایی با مقدار معیار باشند از روش تبدیل داده ها استفاده کرد. تبدیل داده ها با هدف تعدیل متغیرها از جنبه علمی روشی پذیرفته شده است. البته زمانی که اختلاف داده ها با پیش فرض های آماری اندک باشد و به طور تقریبی مفروضات آماری برقرار باشد می توان از تبدیل داده ها صرف نظر کرد.

باید توجه داشت که تبدیل داده ها تا اندازه ای مانند شمشیر دولبه است. حسن این روش این است که می تواند دقت معنی داری تحلیل های آماری را افزایش دهد و عیب آن این است که ممکن است تفسیر داده ها را دشوارتر کند. در نتیجه باید از روش تبدیل داده ها به شیوه ای مدبرانه استفاده کرد.

دشواری در تفسیر داده ها بدین معناست که وقتی داده ها را تبدیل می کنیم، مقدار حداقل و حداکثر و شیوه توزیع متغیر و تمامی شاخص های میانگین و انحراف استاندارد تغییر می کند و با حالت معمول و عادی تفاوت پیدا می کند. مثلا اگر سن افراد که به صورت کمی (نسبی) سنجیده شده است را به توان دو برسانیم شاخص های آماری سن افراد تغییر می کند و با سن های غیر عادی مثل ۲۵۰، ۳۰۰ و غیره مواجه می شویم. یا وقتی متغیری مانند اعتماد اجتماعی داریم و با ۱۰ سوال این متغیر را سنجیدیم و دامنه میانگین این متغیر بین ۱ تا ۵ باشد، لگاریتم گرفتن از این متغیر دامنه نمرات را تغییر می دهد و توضیح و تفسیر متغیر را با مشکل مواجه می کند. یکی از راه های رفع این مشکل این است که هنگام گزارش یافته های توصیفی و شاخص های آماری (مانند میانگین، انحراف استاندارد و مقدار حداقل و حداکثر)؛ یافته ها و شاخص های آماری را هم به صورت عادی (قبل از تبدیل داده ها) و هم بعد از تبدیل داده ها گزارش کنیم.

مرحله چهارم: انتخاب ویژگی ها

در این مرحله به کلمات وزن داده می شود. کلمات با توجه به وزنی که در این مرحله و با استفاده از معیارهای تعریف شده در مرحله بعدی دریافت می کنند، مشخص می شوند. نام دیگر این مرحله استخراج ویژگی است. یعنی در این مرحله ویژگی های موردنظر را مشخص می کنیم تا در مرحله بعدی انتخاب شوند.

در مرحله انتخاب ویژگی، ویژگی هایی که معیارهای تعریف شده را ماکزیمم می کنند، انتخاب خواهند شد. این روش با کاهش نمونه ها سعی در ایجاد یک دسته بندی مناسب دارد.

-تشخیص توییتهای مورد علاقه کاربر

اولا بایستی مرتبط بودن توثیت را با علایق کاربر پیدا کرد. به منظور تعیین میزان ارتباط یک توثیت با موضوعات مورد علاقه کاربر در اغلب کارهای انجام شده از روش TF-IDF و کسینوس زاویه بین بردار کلمات tw_i و pu_i استفاده می گردد. این معیار تعداد تکرار ویژگی ها را در متون بررسی می کند. در این معیار یک میزان آستانه تعریف می شود. آن ویژگی هایی که تعداد تکرارشان بیشتر یا کمتر از میزان آستانه است، حذف می شوند. تکرار زیاد یک ویژگی در اینجا احتمال انتخاب آن ویژگی را بیشتر می کند. از جمله ویژگی های این روش، مقیاس پذیری، سادگی و تأثیر آن است. معیار دیگر مورد استفاده "بهره اطلاعاتی" می باشد. این روش با استفاده از آنتروپی قابل محاسبه است. و آن ویژگی هایی که میزان gain بیشتری دارند را انتخاب می کند. برای افزایش دقت روش پیشنهادی از آنتولوژی استفاده می گردد. آنتولوژی برای مدل سازی شرایط در یک دامنه مورد علاقه و همچنین روابط میان این شرایط استفاده می شود. مهم ترین بخش آنتولوژی نقش کلیدی آن در توسعه وب معنایی است. تحلیل احساسات با استفاده از آنتولوژی به این صورت است که از آنتولوژی جهت استخراج مفاهیم مرتبط استفاده می شود. این بخش به بخش استخراج ویژگی و ویژگی اعمال شده است (در شرح مسئله، بخش استخراج ویژگی ها توضیح داده شده است). در واقع آنتولوژی یک نوع معناشناسی انجام می دهد. کلمات یکسان، ممکن است معانی مختلف و کلمات مختلف، ممکن است معانی یکسان داشته باشند. آنتولوژی کلماتی را که از نظر مفهوم به آن ویژگی ها نزدیک ترند، مشخص می کند.

-بررسی تعداد لایک، تعداد کامنت ها و تعداد ذکر شدن ها:

لایک شدن یک پست توسط کاربر نشان دهنده آن است که کاربر نسبت به آن موضوع حساسیت بیشتری دارد. بنابراین تعداد لایک می تواند پارامتر مهمی در تشخیص احساس کاربر باشد. از طرفی اگر تعداد منشن ها زیاد باشد یعنی کاربران تمایل دارند موضوع را با سایر دوستان خود به اشتراک بگذارند. و همین طور تعداد کامنت ها نشان دهنده آن است که موضوع برای کاربر جذاب بوده و به ارائه ایده خود پرداخته است.

مرحله پنجم: کلاس‌بندی

جدول (۲): معیار ارزیابی

		نتایج پیش‌بینی	
		درست	غلط
		positive	negative
$\begin{matrix} \text{دسته} \\ \text{واقعی} \end{matrix}$	درست	True Positive	True Negative
	غلط	False Positive	False Negative
	positive		
	negative		

شبکه‌های عصبی پیچشی رده‌ای از شبکه‌های عصبی عمیق هستند که معمولاً برای انجام تحلیل‌های تصویری یا گفتاری در یادگیری ماشین استفاده می‌شوند. این شبکه یک الگوریتم یادگیری عمیق است که تصویر ورودی را دریافت می‌کند و به هر یک از اشیا / جنبه‌های موجود در تصویر میزان اهمیت (وزن‌های قابل یادگیری و بایاس) تخصیص می‌دهد و قادر به متمایزسازی آن‌ها از یکدیگر است. در الگوریتم ConvNet مقایسه با دیگر الگوریتم‌های دسته‌بندی به پیش‌پردازش کمتری نیاز است. در حالی که فیلترهای روش‌های اولیه به صورت دستی مهندسی شده‌اند، شبکه عصبی پیچشی، با آموزش دیدن به اندازه کافی، توانایی فراگیری این فیلترها / مشخصات را کسب می‌کند. معماری ConvNet مشابه با الگوی اتصال نورون‌ها در مغز انسان است و از سازمان‌دهی قشر بصری در مغز الهام گرفته شده است. هر نورون به محرک‌ها تنها در منطقه محدودی از میدان بصری که تحت عنوان میدان تاثیر شناخته شده است پاسخ می‌دهد. یک مجموعه از این میدان‌ها برای پوشش دادن کل ناحیه بصری با یکدیگر هم‌پوشانی دارند. ConvNet قادر است به طور موفق و وابستگی‌های زمانی و فضایی را در یک تصویر با استفاده از فیلترهای مرتبط ثبت کند و همچنین، معماری فیلترگذاری بهتری را روی مجموعه داده تصویر به دلیل کاهش تعداد پارامترهای درگیر و استفاده مجدد از وزن‌ها انجام می‌دهد.

۴- ارزیابی

در این بخش به ارزیابی روش پیشنهادی که برحسب مجموعه دادگان تویتر است می‌پردازیم.

۴-۱ مجموعه دادگان

برای ارزیابی روش پیشنهادی بر طبق مراحل زیر عمل کرده و خروجی روش پیشنهادی با سایر الگوریتم‌های یادگیری ماشین مقایسه خواهد شد. برای جمع‌آوری داده در تویتر از روش جریان تویتر استفاده شده است. بر طبق این ابزار توییت‌های مورد نظر بر حسب موضوعات مختلف

جدول (۱): جزئیات مجموعه داده

شماره	مجموعه دادگان	تعداد پست‌ها	تعداد کاربران	تعداد بازنشرها
مجموعه دادگان یک	کوید ۱۹	۲۳۴۵۷	۶۵۸۹۰	۶۵۴۷۸۰
مجموعه دادگان دو	آتشسوزی در جنگل‌های استرالیا	۴۵۰۰	۷۸۹۰۴	۵۶۷۸۹۰
مجموعه دادگان سه	بهار عرب	۵۴۹۸	۴۵۶۷۸۹	۳۷۸۹۰

جمع‌آوری شده است. برای این مقاله از سه موضوع مختلف، کوید ۱۹، آتشسوزی در جنگل‌های استرالیا و بهار عرب استفاده شده است.

دانشی که در مرحله یادگیری مدل تولید می‌شود، می‌بایست در مرحله ارزیابی مورد تحلیل قرار گیرد تا بتوان ارزش آن را تعیین نمود و در پی آن کارایی الگوریتم یادگیرنده مدل را نیز مشخص کرد. این معیارها را می‌توان هم برای مجموعه داده‌های آموزشی در مرحله یادگیری و هم برای مجموعه رکوردهای آزمایشی در مرحله ارزیابی محاسبه نمود.

بالای روش **SASM** به ماهیت این روش برمی‌گردد. در میان روش‌های دیگر، روش **SVM** عملکرد مناسبی داشته است. علت این برتری در قدرت طبقه‌بندی این روش می‌باشد.

در جدول ۴ مقادیر دقیق حساسیت روش‌های مختلف ارزیابی شده را نشان می‌دهد. در بررسی نتایج، نکته حائز اهمیت دقت بالای روش **SASM** پیشنهاد شده در مقایسه با دیگر روش‌هاست. **SASM** در برابر روشی مانند **KNN** بسیار کند است، با این وجود، دقت بالای این روش علی‌رغم کندی آن می‌تواند قابلیت‌های این روش را نشان دهد. دقت

جدول (۴): مقایسه نرخ حساسیت روش‌های پیشنهادی با بقیه روش‌ها

معیار حساسیت	میانگین	مجموعه دادگان اول	مجموعه دادگان دوم	مجموعه دادگان سوم
SASM	۰/۸۵۴۸۱۵۷۷۱	۰/۸۹۶۶۶۶۷۰۰	۰/۸۸۱۲۵۰۰۰	۰/۷۸۶۵۳۰۶۱۲
KNN	۰/۴۵۲۴۰۲۴۲۸۰	۰/۴۱۶۳۶۳۶۰۰	۰/۴۵۶۷۸۵۷۱۴	۰/۴۸۴۰۵۷۹۷۱
نیویزین	۰/۵۳۳۲۷۰۹۷۰	۰/۵۱۶۴۸۳۵۱۶	۰/۵۴۴۲۵۵۳۱۹	۰/۵۳۹۰۷۴۰۷۴
SVM	۰/۷۲۶۳۳۰۹۹۰	۰/۶۹۶۳۶۳۶۰۰	۰/۷۷۸۵۷۱۴۰۰	۰/۷۰۴۵۰۵۷۹۷۱

تشخیص نیز، روش **SVM** به صورت میانگین، بهترین عملکرد را بعد از **SASM** به خود اختصاص داده است. روش **SASM** توانسته است، موارد عدم بازنشر را بهتر از روش‌های دیگر، تشخیص دهد. در این بخش، روش **KNN** بدترین عملکرد را به خود اختصاص داده است.

در جدول ۵، روش‌های پیشنهادی با سه روش دیگر از نظر معیار تشخیص مقایسه شده است. در این ارزیابی، میزان پیش‌بینی درست تصمیمات عدم بازنشر، در معیار تشخیص بسیار موثر است. همان‌طور که در جدول ۵، مشاهده می‌شود، **SASM** بهترین عملکرد را داشته است. در معیار

جدول (۵): مقایسه معیار تشخیص روش‌های پیشنهادی با بقیه روش‌ها

معیار تشخیص	میانگین	مجموعه دادگان اول	مجموعه دادگان دوم	مجموعه دادگان سوم
SASM	۰/۷۳۸۱۶۸۷۹۰	۰/۶۶۰۷۳۱۷۰	۰/۷۸۹۶۹۷۰۰	۰/۷۶۴۰۷۷۶۷
KNN	۰/۳۹۴۱۴۸۵۲۷	۰/۳۵۹۳۴۰۷۰	۰/۳۹۵۳۲۷۱۰	۰/۴۲۷۷۷۷۷۸
نیویزین	۰/۵۴۹۶۹۱۵۴۳	۰/۵۱۰۳۱۲۵۰	۰/۵۶۵۵۹۱۳۹۸	۰/۵۷۳۱۷۰۷۳
SVM	۰/۵۹۱۶۶۸۵۵۷	۰/۴۵۹۱۸۳۷۰	۰/۵۸۵۷۱۴۳۰	۰/۷۳۰۱۵۸۷۳

در ادامه ارزیابی روش‌های پیشنهادی برای پیش‌بینی تصمیم بازنشر کاربر در مواجهه با یک پست، به بررسی معیار نرخ خطای روش‌های پیشنهادی در مقایسه با روش‌های دیگر می‌پردازیم. معیار دقت، یکی از پرکاربردترین معیارها برای ارزیابی روش‌های پیش‌بینی و کلاس‌بندی است. هر چه نسبت میزان پیش‌بینی درست تصمیم بازنشر، به پیش‌بینی‌های نادرست، بیشتر باشد، دقت روش مربوطه بالاتر خواهد بود. جزئیاتی بیشتری از نتایج مقایسه روش‌های پیشنهادی با دیگر روش‌ها از نظر معیار دقت، در جدول ۶ آورده شده است.

جدول (۶): مقایسه نرخ دقت روش‌های پیشنهادی با بقیه روش‌ها

میانگین	مجموعه دادگان اول	مجموعه دادگان دوم	مجموعه دادگان سوم	معیار دقت
۰/۷۶۶۶۷۵۷۲۶	۰/۶۸۵۸۵۷۱۰	۰/۸۶۱۵۳۸۵	۰/۷۵۲۶۳۱۵۷۹	SASM
۰/۴۰۶۳۶۶۱۲۷	۰/۳۴۳۶۷۸۲۰	۰/۴۸۵۷۱۴۳۰	۰/۳۸۹۷۰۵۸۸۲	KNN
۰/۶۴۵۳۸۲۱۷۳	۰/۶۸۲۱۲۱۲۱۲	۰/۶۹۱۷۵۲۵۷۷	۰/۵۶۲۲۷۲۷۳۰	نیویزین
۰/۷۰۷۷۳۰۴۲۶	۰/۵۹۵۲۳۸۱۰	۰/۷۸۸۷۳۲۴۰	۰/۷۳۹۲۲۰۷۷۹	SVM

در جدول ۷، معیار F-measure روش‌های مختلف مقایسه شده است. همان‌طور که پیش‌تر اشاره شد، معیار F-measure بکارگرفته شده در این رساله، میانگین موزون دقت و بازیابی (حساسیت) می‌باشد ($\beta=1$). همان‌طور که اشاره شد، میزان این معیار، ارتباط مستقیمی به دقت و بازیابی روش‌ها دارد. هر دو روش پیشنهادی، بهترین عملکرد را داشته‌اند. نکته جالب در این جدول، پایین بودن مقدار F-measure برای روش KNN است. با وجود دقت نسبتاً خوب این روش، با توجه به ضعف شدید این روش در بازیابی، بدترین عملکرد را در این معیار دارد.

همان‌طور که در جدول ۶ مشاهده می‌شود، روش SASM در این معیار، بهترین عملکرد را داشته است. بعد از این روش ماشین بردار پشتیبان بهترین دقت را در بین روش‌های دیگر داشته است. یکی از دلایل برتری روش‌های پیشنهادی بر اساس معیار دقت در ماهیت این روش‌ها می‌باشد. در اجتماعات برخط ممکن است کاربران دچار تغییر سلیق شوند. این برتری روش پیشنهادی در دقت پیش‌بینی‌های انجام شده بسیار موثر است.

جدول (۷): مقایسه معیار F measure روش‌های پیشنهادی با بقیه روش‌ها

میانگین	مجموعه دادگان اول	مجموعه دادگان دوم	مجموعه دادگان سوم	معیار F measure
۲/۳۸۷۷۸۰۶۱۲	۲/۳	۲/۲۸۳۷۷۵	۲/۵۷۹۵۹۱۸۳۷	SASM
۱/۵۱۳۳۸۴۱۹۷	۱/۶۹۹۰۹۰۹	۱/۱۸۵۷۱۴۳	۱/۶۵۵۳۴۷۳۹۱	KNN
۱/۵۳۱۶۱۲۱۵۲	۱/۵۹۷۰۱۱۲۹۹	۱/۴۴۱۱۲۴۱۲۷	۱/۵۵۶۷۰۱۰۳۰	نیویزین
۱/۷۸۵۸۹۴۵۱۷	۱/۷۱۴۲۸۵۷	۱/۴۴۸۲۷۵۹	۲/۱۹۵۱۲۱۹۵۱	SVM

منفی با استفاده از روش شبکه عصبی چرخشی خواهیم پرداخت. این روش با سه روش دیگر نیز برای ارزیابی مقایسه شد و از سه روش دیگر عملکرد بهتری داشت.

رفتارشناسی کاربران در شبکه‌های اجتماعی یکی از جذاب‌ترین بحث‌های حوزه فناوری اطلاعات در دهه اخیر می‌باشد. رفتارشناسی کاربران این امکان را در اختیار توسعه‌دهندگان فناوری اطلاعات فراهم می‌آورد که با استفاده از نیازسنجی تعاملات و برهم‌کنش علایق و خصوصیات کاربران، به ارائه سرویس بپردازند. این سرویس‌ها ممکن است در قالب ارائه آگهی تبلیغاتی، یا یک فرآیند مدیریت دانش و یا حتی به صورت یک توصیه صورت پذیرد. به عنوان پیشنهاد برای کارهای آینده، می‌توان روی رفتارشناسی انتشار در شبکه‌های اجتماعی مطالعات تکمیلی صورت پذیرد. کاربران تاثیرگذار بر روی سایر کاربران اکتشاف شوند، این کاربران به گونه‌ای هستند که بیشترین احساس مثبت در جذب کاربران برای خرید یک محصول به سایر کاربران می‌دهند و همین‌طور می‌توانند بیشترین احساس منفی را برای دفع کاربران برای خرید یک محصول در شبکه‌های اجتماعی داشته باشند. پست‌های تاثیرگذاری که بیشترین احساس مثبت و منفی را برای جذب یا دفع به خرید یک محصول می‌انجامد را اکتشاف کرد. با اتخاذ روش‌های دیگر برای مدل پیشنهادی

۵- نتیجه‌گیری

تحلیل انتشار اطلاعات و نفوذ اجتماعی در شبکه‌های اجتماعی دارای کاربردهای بسیار زیادی در جهان واقعی دارد. یکی از مثال‌های کاربردی آن پیشنهادی نفوذ در بازاریابی و ویروسی می‌باشد. تعیین کاربران تأثیر گذار به عنوان یکی از اصلی‌ترین موضوعات موجود در شبکه‌های اجتماعی می‌باشد که اهمیت فراوانی دارد. چنانچه این کاربران به صورت دقیق‌تری شناسایی شوند، عملیاتی که بر مبنای این کاربران انجام می‌شود با نفوذتر خواهد بود. هدف از انجام این پروژه ارائه روشی برای تحلیل احساسات مشتری در رسانه‌های اجتماعی جهت استفاده در سامانه‌های تبلیغات است. برای این امر از API جست و جوی تویتر استفاده شده است. این مجموعه داده شامل اطلاعات مرتبط با کاربران و فعالیت‌هایشان در بستر شبکه اجتماعی تویتر بوده است که در پایگاه داده ذخیره شده است. این داده‌های ذخیره شده برای آماده سازی به سه دسته ساختار یافته، نیمه ساختار یافته و غیر ساختار یافته تقسیم شده است و برای هر کدام از ساختارها مراحل برای آماده سازی داده انجام شده است. سپس داده‌ها به فرمت مناسب تبدیل شده است و ویژگی‌های مورد نظر استخراج می‌شود و در نهایت به طبقه‌بندی کاربران براساس احساسات مثبت و

- Convolutional Neural Networks.** in Proc of the ۲۲nd ACM International Conf. on Multimedia, pp. ۹۷-۱۰۶, Orland, FL, USA, ۱۸-۱۹.
- [۱۶] Zhang, Y., Chen, M., Liu, L., Wang, Y. (۲۰۱۷). **An Effective Convolutional Neural Network Model for Chinese Sentiment Analysis.** in Proc AIP Conf. Proc., vol. ۱۸۳۶, pp. ۰۲۰۰۸۴, Rome, Italy, ۲۷-۲۹.
- [۱۷] on Empirical Methods in Natural Language Processing, EMNLP'۱۳, vol. ۱۶۳۱, pp. ۱۶۳۱-۱۶۴۲, Seattle, WA, USA, ۱۸-۲۱.

می‌توان آن را مقیاس پذیرتر نمود. با دخیل کردن ویژگی‌هایی همچون فرهنگ، نژاد، قومیت به مدل پیشنهادی می‌توان نتایج گسترده‌تری را کسب نمود. با ترکیب مدل‌های طبقه‌بندی و قراردادن وزن‌های متفاوت به ویژگی‌های تأثیرگذار، عملکرد روش پیشنهادی را افزایش داد. از طرفی علائق کاربران با گذشت زمان تغییر می‌کند، این تغییرات یا بستگی به برهه‌های خاص زمانی دارد، مانند اوایل سال جدید و یا بر اثر تغییر طبع کاربر با گذر زمان ایجاد می‌شود، با در نظر گرفتن این پویایی در شبکه‌های اجتماعی می‌توان روش پیشنهادی را انعطاف‌پذیرتر نمود.

منابع و ماخذ

- [۱] Arora, L., Singh, P., Bhatt, V., Sharma, B. (۲۰۲۱). **Understanding and Managing Customer Engagement through Social Customer Relationship Management.** Journal of Decision Systems, ۱-۲۱.
- [۲] Bagheri A., Saraei, M. (۲۰۱۴). **Persian Sentiment Analyzer: a Framework Based on a Novel Feature Selection Method.** International Journal of Artificial Intelligence™, Vol. ۱۲, No. ۲, pp. ۱۱۵-۱۲۹.
- [۳] Chen, Z S., Zhang, X., Govindan, K., Wang, X. J., Chin, K. S. (۲۰۲۱). **Third-Party Reverses Logistics Provider Selection: A Computational Semantic Analysis-Based Multi-Perspective Multi-Attribute Decision-Making Approach.** Expert Systems with Applications, ۱۶۶, ۱۱۴۰۵۱.
- [۴] Cieliebak, M., Deriu, J., Egger, D., Uzdilli, F. (۲۰۱۷). **A Twitter Corpus and Benchmark Resources for German Sentiment Analysis.** in Proc of the ۵th Ine, Workshop on Natural Language Processing for Social Media, SocialNLP, pp. ۴۵-۵۱, Boston, USA.
- [۵] Collobert, R., Weston, J., Bottou, L., Karlen, M., Kavukcuoglu, K., Kuksa, P. (۲۰۱۱). **Natural Language Processing (Almost) from Scratch.** Journal of Machine Learning Research, vol. ۱۲, No. ۷۶, pp. ۲۴۹۳-۲۵۳۷.
- [۶] Cortes C., Vapnik, V. (۱۹۹۵). **Support-Vector Networks.** Machine Learning, Vol. ۲۰, No. ۳, pp. ۲۷۳-۲۹۷.
- [۷] Dos Santos C.N., Gatti, M. (۲۰۱۴). **Deep Convolutional Neural Networks for Sentiment Analysis of Short Texts.** in Proc of the ۲۵th International Conf. on Computational Linguistics, COLING'۱۴, pp. ۶۹-۷۸, Dublin, Ireland, ۲۵-۲۹.
- [۸] Jaynes, E.T. (۱۹۵۷). **Information Theory and Statistical Mechanics.** Physical Review, Vol. ۱۰۶, No. ۴, pp. ۶۲۰.
- [۹] Maulud, D. H. (۲۰۲۱). **State of Art for Semantic Analysis of Natural Language Processing.** Qubahan Academic Journal ۱,۲, ۲۱-۲۸.
- [۱۰] Mikolov, T., I. Sutskever, I., Chen, K., Corrado, G.S., Dean, J. (۲۰۱۳). **Distributed Representations of Words and Phrases and their Compositionality.** in Proc Advances in Neural Information Processing Systems, NIPS'۱۳, pp. ۳۱۱۱-۳۱۱۹, Lake Tahoe, CA, USA, ۵-۱۰.
- [۱۱] Neethu M.S., Rajasree, R. (۲۰۱۳). **Sentiment Analysis in Twitter using Machine Learning Techniques.** in Proc IEEE ۴th Int. Conf. on, Computing, Communications and Networking Technologies, ICCCNT'۱۳, ۵ pp., Tiruchengode, India, ۴-۶.
- [۱۲] Roshanfekar, B., Khadivi, S., Rahmati, M. (۲۰۱۷). **Sentiment Analysis using Deep Learning on Persian Texts.** in Iranian Conf, on Electrical Engineering, ICEE'۱۷, pp. ۱۵۰۳-۱۵۰۸, Tehran, Iran, ۲-۴.
- [۱۳] Shearer, E., Amy, M. (۲۰۲۱). **News Use across Social Media Platforms in ۲۰۲۰.**
- [۱۴] Socher, R., Perelygin, A., Wu, J.Y., Chuang, J., Manning, C.D., Ng, A.Y., Potts, C. (۲۰۱۳). **Recursive Deep Models for Semantic Compositionality over a Sentiment Treebank.** in Proc of the Conf.
- [۱۵] Wang, K., Wang, X., Lin, L., Wang, M., Zuo, W. (۲۰۱۴). **۳D human Activity Recognition with Reconfigurable**