



چارچوب ترکیبی مبتنی بر NLP و یادگیری عمیق برای تحلیل احساسات کاربران در اینترنت اشیاء اجتماعی

پریسا دلفانی^۱، پرستو بهی^۲، مهشید علیان^۳، مائده فلزی نصیری^۴

۱- گروه مهندسی کامپیوتر، دانشگاه ملی مهارت، تهران، ایران Parisa.delfani.16768@gmail.com

۲- گروه مهندسی کامپیوتر، دانشگاه ملی مهارت، تهران، ایران Parastoobehi@gmail.com

۳- گروه مهندسی کامپیوتر، دانشگاه ملی مهارت، تهران، ایران Mahshidalian1383@gmail.com

۴- گروه مهندسی کامپیوتر، دانشگاه ملی مهارت، تهران، ایران maedenasiri84@gmail.com

چکیده

این پژوهش به طراحی و ارائه یک چارچوب ترکیبی از مدل های یادگیری ماشین و یادگیری عمیق، برای تحلیل احساسات و درک بهتر نظرات کاربران در بستر اینترنت اشیاء اجتماعی می پردازد. در طراحی این چارچوب، مطالعات پیشین نیز مورد بررسی قرار گرفتند؛ از جمله پژوهشی بر پلتفرم Tokopedia که در آن ۵۰۰۰ نظر کاربران با استفاده از الگوریتم های SVM و Naive Bayes تحلیل شده و مطالعات دیگری در مورد داده های استخراج شده از توییتر با مجموع ۱۱۹۹۷ توییت و داده های گرفته شده از آمازون با مجموع ۱۰۰۰ داده، که با استفاده از SVM و Random Forest و ترکیب آن دو تحلیل شده اند. همچنین به مدل های عمیق از جمله CNN ها، RNN ها و مدل های مبتنی بر ترنسفورمر می پردازیم. با توجه به مقایسه های انجام شده در مطالعات پیشین، الگوریتم ترکیبی RFSVM با روش وزن دهی TF-IDF و مدل های مبتنی بر ترنسفورمر در مقایسه با مدل ها و الگوریتم های دیگر عملکرد بهتری داشته اند. در نهایت، چارچوبی ترکیبی از RFSVM و مدل های مبتنی بر ترنسفورمر BERT و GPT پیشنهاد میشود که عملکرد بهتری در تجزیه جزئیات مربوط به احساسات کاربران در IOT اجتماعی دارند.

کلید واژه: تحلیل احساسات، اینترنت اشیاء، یادگیری ماشین، پردازش زبان طبیعی، تحلیل بلادرنگ احساسات، شهر هوشمند

۱. مقدمه:

در این مقاله به چگونگی تجزیه تحلیل نظرات و احساسات کاربران، استخراج اطلاعات مفید و استفاده از آنها در زمینه های مختلف پرداخته می شود. در این زمینه ما با مسائلی همچون مشکلات مربوط به حفظ حریم خصوصی، داده های تکراری یا نویز دار، غلط های املائی، متون کوتاه و غیررسمی، تشخیص کنایه ها و مواردی که در ادامه به آنها اشاره خواهد شد مواجه هستیم. در این مقاله، نقش فناوری های زبان طبیعی و هوش مصنوعی در استخراج اطلاعات احساسی از نظرات کاربران مورد بررسی قرار می گیرد. همچنین، یک چارچوب ترکیبی برای تحلیل احساسات در بستر داده های اجتماعی اینترنت اشیاء ارائه خواهد شد که با روش های پیشرفته ی NLP و الگوریتم های یادگیری ماشین، بتواند احساسات کاربران را پردازش و طبقه بندی کند.



با افزایش داده‌ها در بستر اینترنت اشیاء، تحلیل احساسات به ابزاری کلیدی در خدمات هوشمند تبدیل شده است. ترکیب NLP و یادگیری ماشین، امکان درک عمیق‌تر نیازهای کاربران را فراهم می‌سازد و در حوزه‌هایی مانند شهر هوشمند و خدمات مشتری اهمیت دارد.

در بخش دوم، علاوه بر مرور کلی بر چالش‌ها، فرصت‌ها و خلاءهای پژوهشی، توضیحات جامعی در رابطه با تحلیل احساسات و پردازش زبان طبیعی ارائه داده ایم و در نهایت به مدل‌های هوش مصنوعی در تحلیل احساسات و کاربرد آنها در IOT پرداخته ایم. در بخش سوم کاربرد های تحلیل احساسات در شهر هوشمند، اینترنت اشیا، خدمات مشتری هوشمند و همچنین اپلیکیشن‌های موبایل را مورد بررسی قرار خواهیم داد. در بخش چهارم در چهارچوب پیشنهادی ارائه شده، مراحل فرآیند تحلیل احساسات را از جمع‌آوری داده تا ارزیابی مدل بررسی خواهیم کرد. بخش پنجم مقایسه‌ی مدل‌های یادگیری ماشین با داده‌های واقعی از پلتفرم‌های توئیتر، آمازون و Tokopedia و همچنین مقایسه الگوریتم‌های یادگیری عمیق با یکدیگر می‌باشد. در نهایت، بخش ششم به نتیجه‌گیری، تحلیل نتایج و پیشنهادهایی برای پژوهش‌های آینده اختصاص دارد.

۲. مرور ادبیات

۲/۱. چالش‌ها و فرصت‌ها

فرصت‌ها:

یکی از فرصت‌های مهم، پیشرفت تحلیل‌های چندوجهی است، جایی که داده‌های صوتی، تصویری و متنی به‌صورت هم‌زمان برای درک عمیق‌تر احساسات کاربران ترکیب می‌شوند. این رویکرد می‌تواند تا حدود 10.5% به کاهش خطا و افزایش دقت مدل‌ها منجر شود.[3]

فرصت بعدی استفاده از یادگیری فدرال (Federated Learning) است. این روش به سیستم‌ها امکان می‌دهد بدون انتقال داده‌های خام و با حفظ حریم خصوصی، مدل‌های تحلیلی را به‌صورت توزیع‌شده آموزش دهند. در این چارچوب، داده‌های کاربران روی دستگاه‌های خود باقی می‌مانند و تنها پارامترهای به‌روزرسانی‌شده‌ی مدل‌ها به اشتراک گذاشته می‌شوند.[4]

همچنین رایانش احساسی، حوزه‌ای نوظهور در این زمینه است که هدف آن توسعه‌ی سامانه‌هایی است که قادرند از طریق داده‌هایی مانند سیگنال‌های فیزیولوژیکی، صوت یا تصویر؛ احساسات کاربران را شناسایی و به آن‌ها واکنش مناسب نشان دهند.[3]

فرصت دیگر، سامانه‌هایی هستند که بتوانند به صورت بلادرنگ به محتوای احساسی تولیدشده در شبکه‌های اجتماعی واکنش نشان دهند.[6]



چالش ها :

تحلیل احساسات شامل فرآیند شناسایی نگرش کاربران است که دارای بار احساسی مثبت، منفی یا خنثی است. اولین چالش، وجود همین جملات خنثی است که مدل‌های تحلیل احساسات توانایی تشخیص دقیق آن‌ها را ندارند. [1]

یکی دیگر از چالش‌ها، شناسایی طعنه و کنایه در متن است که با نشانگرهای زبانی مانند بیان‌های اغراق‌آمیز و علائم نگارشی شناسایی می‌شوند. چون مدل‌های سنتی در این زمینه ناتوان هستند، چندین مجموعه داده‌ی حاشیه‌نویسی شده برای تسهیل تحقیقات در حوزه طعنه، کنایه و ابهام ایجاد شده‌اند؛ اما حاشیه‌نویسی نیز چالش‌های مختص به خود را دارد. [9]

از طرفی در تحلیل احساسات مبتنی بر شبکه‌های اجتماعی، وجود نویز، غلط‌های املایی، داده‌های تکراری و اسپم از جمله مشکلاتی است که دقت مدل‌ها را تحت تأثیر قرار می‌دهد. [3]. یکی از عوامل ایجاد نویز، تغییر زبان (کدسوئیچینگ) است که در آن چند زبان در یک جمله استفاده می‌شود که روندهای سنتی NLP را مختل می‌کند. [9]

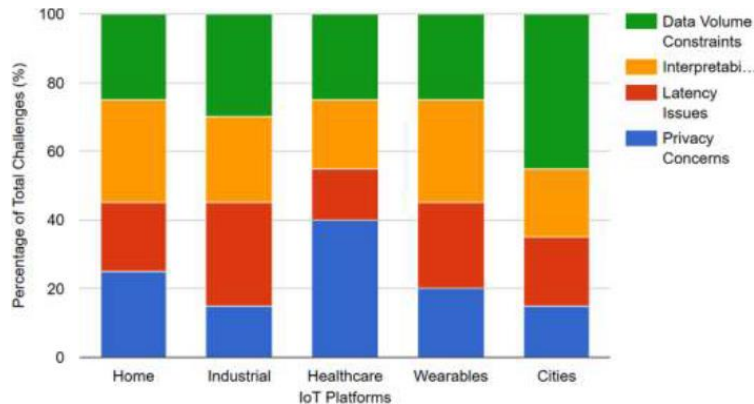
چالش‌های مربوط به حفظ حریم خصوصی نیز اهمیت دارند. برای مثال در صورت عدم حفاظت از داده‌ها، امکان نقض حریم خصوصی، جریمه‌های قانونی و از بین رفتن اعتماد کاربران وجود دارد. [3] سازمان‌ها در راستای مقرراتی مانند GDPR و CCPA موظفند تنه‌ل داده‌های ضروری را جمع‌آوری و آن‌ها را به‌صورت ایمن ذخیره کنند. [4]

کمبود ابزارهای پردازش زبان طبیعی برای زبان‌های غیرانگلیسی، چالش‌های مهمی را در تحلیل احساسات چندزبانه ایجاد می‌کند. [3] بسیاری از ابزارها و فرهنگ لغت تحلیل احساسات عمدتاً برای زبان انگلیسی توسعه یافته‌اند که منجر به کاهش عملکرد و قابلیت انتقال به زبان‌های دیگر می‌شود. همچنین متون غیرانگلیسی اغلب واژگان غیررسمی هستند که در داده‌های آموزشی استاندارد وجود ندارند و باعث ناهماهنگی در تشخیص احساسات بین زبان‌ها می‌شود. [9]

لازم به ذکر است روش‌های سنتی غالباً در مواجهه با چندمعنایی بودن زبان غیررسمی شبکه‌های اجتماعی و اصطلاحات اینترنتی یا مخفف‌ها، دچار مشکل می‌شوند و مدل‌هایی که روی منابع رسمی مانند مقالات خبری یا ویکی‌پدیا آموزش دیده‌اند، اغلب در داده‌های غیررسمی عملکرد ضعیفی دارند. بنابراین، عبارات اصطلاحی و محاوره‌ای نیازمند مدل‌های معنایی آموزش‌دیده روی پیکره‌های غیررسمی هستند. [9]

در حوزه‌ی اینترنت اشیا نیز تحلیل احساسات با چالش‌هایی روبه‌روست؛ در شبکه‌های گسترده‌ی IoT، دستگاه‌ها به صورت مداوم داده تولید می‌کنند و این داده‌ها باید به صورت بلادرنگ تحلیل شوند تا بتوان از آنها استفاده کرد. الگوریتم‌های سنتی تحلیل احساسات برای این کار کند هستند و نمیتوانند به موقع نتیجه دهند. [6] علاوه بر این، چندوجهی بودن داده‌ها موجب پیچیدگی در طراحی الگوریتم‌های تحلیل احساسات در محیط‌های اجتماعی مبتنی بر IoT می‌شود. [3]

در شکل ۱، توزیع چالش‌های اصلی در پنج پلتفرم رایج IoT شامل خانه‌های هوشمند، صنعت، سلامت، پوشیدنی‌ها و شهرهای هوشمند نشان داده شده است. این نمودار سهم نسبی مشکلاتی چون نگرانی‌های حریم خصوصی، تأخیر در پاسخ‌گویی، سختی تفسیر مدل‌ها و محدودیت‌های حجم داده را در هر حوزه مشخص می‌سازد. [6]



شکل ۱ - توزیع درصدی چالش‌های اصلی در پلتفرم‌های مختلف IoT (نگرانی‌های حریم خصوصی، تأخیر، تفسیرپذیری، و حجم داده). [6]

۲/۲. خلاء پژوهش

با وجود پیشرفت‌های روزافزون در حوزه‌ی تحلیل احساسات، همچنان خلأهای قابل توجهی در این زمینه وجود دارد. بخش عمده‌ای از پژوهش‌های انجام‌شده، تمرکز خود را بر تحلیل داده‌های متنی ساختاریافته، نظیر نظرات کاربران در شبکه‌های اجتماعی یا پلتفرم‌های فروش محدود کرده‌اند، در حالی که تحلیل احساسات در داده‌های غیرساختاریافته یا رفتاری، به‌ویژه داده‌های تولیدشده توسط دستگاه‌های IoT، تا حد زیادی نادیده گرفته شده است. علاوه بر این در بسیاری از پژوهش‌ها تحلیل احساسات به صورت آفلاین انجام می‌شود و کمتر به تحلیل بلادرنگ در بسترهای پویا مانند شهر هوشمند، خدمات مشتری هوشمند و اپلیکیشن‌های دیجیتال پرداخته شده است. یکی دیگر از خلأهای اساسی، نبود یک چارچوب ترکیبی، ماژولار و بلادرنگ است که بتواند داده‌های متنی، رفتاری و حسگری را به‌صورت یکپارچه و هم‌زمان تحلیل کند. در این پژوهش، تلاش می‌شود این شکاف‌ها شناسایی و با ارائه‌ی چارچوبی پیشنهادی، تا حد امکان برطرف شوند.

۲/۳. پردازش زبان طبیعی (NLP – Natural language processing)

پردازش زبان طبیعی (NLP) ستون فقرات سیستم‌های خدمات مشتری مبتنی بر هوش مصنوعی است که به ماشین‌ها امکان می‌دهد زبان انسانی را پردازش و درک کنند. با بهره‌گیری از NLP، سیستم‌های هوش مصنوعی قادر به انجام گفت‌وگوهای پیچیده و مبتنی بر زمینه (context) هستند، که برای ارتقاء تعاملات مشتری حیاتی است. این سیستم‌ها مزایای متعددی از جمله

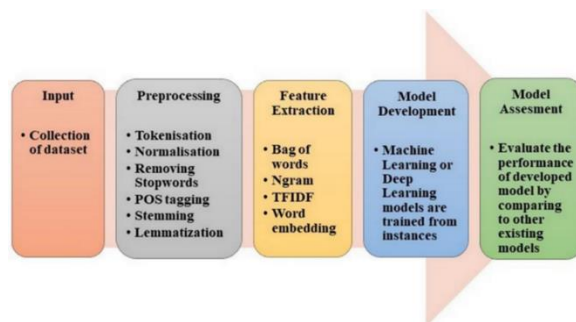
شخصی سازی، کارایی محاسباتی، مقیاس پذیری و دسترسی ۷/۲۴، ارائه می دهند. تکنیک های کلیدی NLP مانند تحلیل احساسات (sentiment analysis)، پایه های پلتفرم های خدمات مشتری مبتنی بر هوش مصنوعی را شکل می دهند و به کسب و کارها کمک می کنند تا تعاملات پیچیده را در مقیاس بالا مدیریت کنند. [4]

پردازش زبان طبیعی از طریق ترکیب لایه های مختلف دانش زبانی این امکان را پیدا میکند تا پرسش های مشتری را به طور دقیق تفسیر کرده و پاسخ های مرتبط تولید کند. این لایه ها عبارتند از:

- تحلیل نحوی (Syntactic Analysis): تحلیل نحوی ساختار گرامری جملات را بررسی می کند. این تحلیل از روش های تجزیه گر مانند تحلیل سازه ای و تحلیل وابستگی، برای برقراری ارتباط میان واژه ها در یک عبارت استفاده می کند. [4]
- تحلیل معناشناختی (Semantic Analysis): تحلیل معناشناختی، سعی میکند معنای دقیق واژه ها و عبارات را در بستر جمله تشخیص دهد. [4]
- تحلیل کاربردشناختی (Pragmatic Analysis): تحلیل کاربردشناختی، بر نیت گوینده و بافت گفتگو تمرکز دارد؛ یعنی بررسی میکند که کاربر با بیان یک جمله دقیقاً چه هدفی دارد و چگونه باید به آن پاسخ داده شود. [4]

۲/۴. تحلیل احساسات (Sentiment Analysis)

تحلیل احساسات که با عنوان استخراج نظر (opinion mining) نیز شناخته می شود، به یکی از روش های کلیدی در استخراج اطلاعات ذهنی از متن تبدیل شده است، به ویژه در رسانه های اجتماعی که کاربران مرتباً احساسات و نظرات را ابراز می کنند. [9] تحلیل احساسات، زیرمجموعه ای از NLP به شمار میرود که هدف آن استخراج و دسته بندی خودکار احساسات موجود در متن، صوت یا داده های دیگر است. [6] این تحلیل معمولاً احساسات را در سه دسته ی مثبت، منفی و خنثی طبقه بندی میکند. [4] تحلیل احساسات با بهره گیری از لایه های نحوی، معنایی و کاربردشناختی، امکان درک عمیق تری از زبان را فراهم میسازد و به سیستم های هوش مصنوعی کمک میکند تا پاسخ هایی متناسب با نیت و زمینه ی کاربر ارائه دهند. مدل های پیشرفته ی هوش مصنوعی ویژگی های معنایی مرتبط با احساس را از متن مشتری استخراج کرده و با استفاده از طبقه بندی یادگیری ماشین نوع احساس را مشخص می کنند. برای مثال، اگر مشتری بازخوردی مانند «من از خدمات شما بسیار ناراضی



هستم» ثبت کند، سیستم می تواند کلمه ی «ناراضی» را به عنوان نشانگر قوی احساس منفی شناسایی کند. [4] شکل ۲ به پنج مرحله ی پایه تحلیل و شناسایی احساسات میپردازد.

شکل ۲- مراحل پایه برای انجام تحلیل احساسات و شناسایی احساسات [9]

فرا تر از داده های متنی، شناسایی احساسات به تحلیل چندوجهی گسترش یافته است که با ادغام داده های بصری، صوتی و متنی، دقت را افزایش می دهد. محتوای شبکه های اجتماعی اغلب شامل تصاویر، ایموجی ها، میم ها و ویدیو هایی است که احساسات متنی را تکمیل کرده و موقعیت های تفسیری پیچیده ای ایجاد می کنند. تحلیل احساسات چندوجهی با ترکیب حالت های اضافی مانند ایموجی ها، هشتگ ها، تصاویر و ویدئوها، شناسایی احساسات متنی را تقویت می کند. [9]

۲/۵. مدل های هوش مصنوعی در تحلیل احساسات مشتری

تحول در مدل ها و چارچوب های هوش مصنوعی، اتوماسیون خدمات مشتری را با فراهم کردن توانایی مدیریت حجم بالای تعاملات با کارایی، شخصی سازی و مقیاس پذیری بیشتر دگرگون کرده است. این پیشرفت ها به سیستم ها امکان می دهد نیازهای مشتری را پیش بینی کنند و وظایف روتین را خودکار سازند. در هسته ی این تحول سه نوع مدل وجود دارد: [4]

• یادگیری ماشین

مدل های یادگیری ماشین در وظایفی مانند دسته بندی پرسش ها، تحلیل احساسات و پیش بینی رفتار مشتری کارآمد هستند. هر مدل اصول یادگیری متمایزی را برای حل مسائل دسته بندی، رگرسیون و پیش بینی در جریان کار خدمات مشتری به کار می گیرد. [4]

مدل های کلیدی یادگیری ماشین که در اتوماسیون خدمات مشتری کاربرد دارند:

* جنگل تصادفی (Random Forest): الگوریتم جنگل تصادفی قادر است حجم زیادی از داده ها را با دقت بالا طبقه بندی کند. این الگوریتم در زمان آموزش، چندین درخت تصمیم ایجاد می کند و کلاسی را به عنوان خروجی ارائه می دهد که بیشترین تعداد دفعات انتخاب شده توسط درخت های منفرد، باشد. [1]

* ماشین بردار پشتیبان (Support Vector Machine - SVM): الگوریتم SVM یک مرز تصمیم گیری بهینه ایجاد کند که داده های مربوط به کلاس های مختلف را با بیشترین فاصله از یکدیگر جدا کند. این الگوریتم به دلیل دقت بالا و مقاومت در برابر بیش برآزش (Overfitting)، به ویژه در داده های متنی، کاربرد گسترده ای در پژوهش های تحلیل احساسات دارد. [2]

* نایو بیز (Naive Bayes): الگوریتم نایو بیز بر اساس قانون بیز و فرض استقلال بین ویژگی ها عمل می کند. این الگوریتم بر اساس احتمال تصمیم می گیرد که یک متن به کدام دسته ی احساسی تعلق دارد: مثبت یا منفی. الگوریتم نایو بیز در کاربردهایی مانند تحلیل احساسات متون کوتاه یا داده های حجیم، به دلیل سادگی و سرعت بالا، بسیار مفید واقع می شود. [2]

• یادگیری عمیق

یادگیری عمیق تحول بزرگی در تحلیل احساسات ایجاد کرده است چرا که امکان یادگیری نمایه‌های سلسله‌مراتبی و متنی را بدون نیاز به مهندسی دستی ویژگی‌ها فراهم می‌کند. معماری‌های یادگیری عمیق برای وظایفی مانند هوش مصنوعی مکالمه‌ای، تحلیل احساسات و حل پرسش‌ها کارآمدند. سه دسته کلیدی مدل‌های یادگیری عمیق عبارتند از:

* شبکه‌های عصبی کانولوشنی (CNN): CNN ها که در ابتدا برای شناسایی تصویر طراحی شده بودند، با موفقیت برای وظایف NLP، به‌ویژه در کشف الگوهای محلی متن مانند n-gram های حاوی احساس به کار می‌روند. [9]

با اعمال فیلترهای کانولوشن بر کلمات، CNN ها قادر به شناسایی وابستگی‌های مکانی و ویژگی‌های موقعیت‌ناپذیر مرتبط با طبقه‌بندی احساس هستند. مطالعات نشان داده‌اند CNN ها در متن‌های کوتاه مانند توئیتهای یا نقدهای محصول که در آنها الگوهای موضعی مانند منفی‌سازی، شدت‌دهنده‌ها یا تعدیل‌کننده‌های احساس نقش اصلی دارند، عملکرد خوبی دارند. [9]

* شبکه‌های عصبی بازگشتی (RNN):

RNN-Recurrent Neural Network: شبکه‌های RNN از اولین مدل‌های یادگیری عمیق بوده‌اند که برای پردازش داده‌های دنباله‌ای طراحی شده‌اند. این مدل‌ها ساختاری دارند که اطلاعات ورودی‌های قبلی را در حافظه داخلی نگه می‌دارند و هر ورودی جدید را در زمینه آن اطلاعات پردازش می‌کنند. [4] RNN ها به ویژه برای تحلیل احساسات مناسب هستند چون توانایی پردازش داده‌های ترتیبی و حفظ اطلاعات در طول گام‌های زمانی را دارند. [9]

با این حال، RNN های استاندارد به مشکل ناپدید شدن گرادیان دچارند و وقتی دنباله‌ها طولانی می‌شوند، مدل توانایی حفظ وابستگی‌های بلندمدت را از دست می‌دهد که محدودیتی برای پردازش مکالمات پیچیده یا چند مرحله‌ای است. [4]

همچنین شبکه‌های بازگشتی دوطرفه (Bi-RNN) پیشنهاد شده‌اند تا هم بافت گذشته و هم آینده را درک کنند و بدین ترتیب دقت طبقه‌بندی احساسات را افزایش دهند. [9]

LSTM-Long Short Term Memory: برای غلبه بر محدودیت‌های RNN ها، شبکه‌های LSTM توسعه یافته‌اند. شبکه‌های حافظه بلندمدت (LSTM) با مکانیسم‌های دروازه‌دار، محدودیت‌های RNN های سنتی را برطرف کرده‌اند و امکان نگهداری و فراموشی انتخابی اطلاعات در توالی‌های طولانی را فراهم می‌کنند. [9] این باعث می‌شود معماری‌های مبتنی بر LSTM در کاربردهای چت‌بات‌ها مؤثر باشند، جایی که حفظ پیوستگی دیالوگ و درک قصد کاربر در چند پیام اهمیت دارد. [4]

GRU- Gated Recurrent Unit: واحدهای بازگشتی دروازه‌دار (GRU) معماری ساده‌تری نسبت به LSTM دارند، به‌طوری که در آنها دروازه فراموشی و ورودی در یک دروازه به نام دروازه به‌روزرسانی ترکیب شده‌اند و در عین حال عملکرد مشابهی در وظایف احساسات ارائه می‌دهند. GRU های دوطرفه (Bi-GRU) با در نظر گرفتن جریان معنایی جلو و عقب، به ویژه در داده‌های چندزبانه یا کدمیکس شده، پیش‌بینی احساسات را بهبود می‌بخشند. [9]

* معماری‌های مبتنی بر ترنسفورمر:

معماری‌های مبتنی بر ترنسفورمر که با مدل‌هایی مانند BERT و GPT معرفی شدند، با استفاده از رویکرد پردازش موازی، انقلابی در پردازش زبان طبیعی (NLP) ایجاد کردند. [4]

BERT: مدلی رمزگذار محور است که متن را به صورت دوطرفه (هم از چپ به راست و هم از راست به چپ) تحلیل می‌کند. این ویژگی باعث درک بهتر زمینه معنایی واژه‌ها می‌شود. کاربرد اصلی BERT شامل دسته‌بندی پرسش‌ها، شناسایی موجودیت‌های نامدار و تحلیل احساسات می‌شود. [4]

GPT: مدل‌های GPT در تولید پاسخ‌های شبیه به انسان تخصص دارند و در تولید متن منسجم با پیش‌بینی هر کلمه بر اساس زمینه‌ی پیشین، بسیار قوی هستند. GPT در برنامه‌های مکالمه‌ای که حفظ جریان طبیعی گفت‌وگو اهمیت دارد، بسیار موثر است و بر خلاف BERT که بیشتر برای درک زبان طراحی شده است، در وظایف تولید متن عملکرد برجسته‌ای دارد. [4]

معماری‌های مبتنی بر ترنسفورمر با بهره‌گیری از مکانیزم‌های خود-توجه و رمزگذاری دوسویه بافت عملکردی در سطح عالی به‌دست آورده‌اند. این مدل‌ها روی پیکره‌های عظیم آموزش داده شده و برای وظایف طبقه‌بندی احساسات تنظیم دقیق (fine-tune) شده‌اند، که منجر به بهبود عملکرد در مجموعه داده‌های مرجع شده است. یادگیری انتقالی و مدل‌های زبانی از پیش آموزش‌دیده، به‌ویژه BERT، RoBERTa، و XLNet، امکان تشخیص دقیق‌تر احساسات را فراهم کرده‌اند. [9]

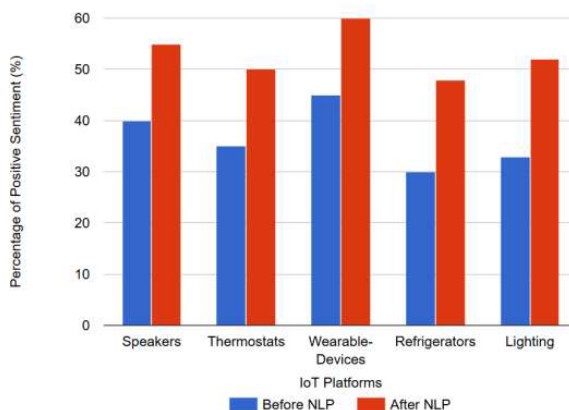
• یادگیری تقویتی (Reinforcement Learning)

یادگیری تقویتی (RL) یک الگوی یادگیری پویا برای سیستم‌های هوش مصنوعی است که به پلتفرم‌های خدمات مشتری این امکان را می‌دهد تا از بازخوردها برای بهبود تدریجی نتایج تعاملات استفاده کنند. این نوع یادگیری، به سیستم اجازه می‌دهد بر اساس پاداش یا جریمه‌ای که از اقداماتش دریافت می‌کند، بهترین روش برای مدیریت یک موقعیت خاص را یاد بگیرد. [4]

۲.۶. تحلیل احساسات در IOT :

داده‌ها، عموماً به دو دسته‌ی ساختاریافته و بدون ساختار تقسیم می‌شوند. در داده‌های ساختار یافته، کاربران مستقیماً نظر خود را گفته‌اند، مانند پست‌های شبکه‌های اجتماعی، نظرات محصول یا پاسخ‌های نظرسنجی؛ تفسیر و دسته‌بندی این نوع داده‌ها ساده است. اما تحلیل داده‌های بدون ساختار مانند اندازه‌گیری‌های حسگرها، سیگنال‌های رفتاری یا لاگ وضعیت دستگاه‌ها که ممکن است نشان دهنده رضایت یا نارضایتی کاربر باشند، پیچیده‌تر است. این نوع داده‌ها ممکن است اطلاعات احساسی در خود داشته باشند، اما مانند بازخوردهای متنی واضح نیستند. به عنوان مثال افزایش دمای تنظیم‌شده روی ترموستات هوشمند ممکن است نشان دهنده ناراحتی کاربر باشد، یا تکرار زیاد یک فرمان صوتی به دستیار هوشمند ممکن است نشان‌دهنده این باشد که کاربر از پاسخ‌دهی دستگاه راضی نیست. این نوع داده‌ها اگرچه مثل متن مستقیم نیستند، ولی اگر به درستی تحلیل بشوند، می‌توانند نتایج مفیدی ارائه بدهند. تحلیل احساسات در ابتدا بر داده‌های متنی ساختار یافته متمرکز بود، اما با پیشرفتهای اخیر در مدل‌های مبتنی بر ترنسفورمر مانند BERT و GPT، این قابلیت‌ها به طور چشمگیری توسعه یافته‌اند و توانایی تحلیل داده

های بدون ساختار IOT را پیدا کرده اند. [6] در شکل ۳ مقایسه ای از درصد احساسات مثبت کاربران نسبت به پلتفرم های مختلف اینترنت اشیا آورده شده.



شکل ۳ - مقایسه ی درصد احساسات مثبت کاربران نسبت به پنج پلتفرم مختلف اینترنت اشیا (شامل اسپیکرهای هوشمند، ترموستات ها، دستگاه های پوشیدنی، یخچال ها و سیستم های روشنایی) پیش از به کارگیری روش های پردازش زبان طبیعی و پس از آن. [6]

روش های NLP برای پردازش داده های پیچیده و چندمنظوره در محیط های IOT کاملاً ضروری هستند. مدل های مدرن NLP، به ویژه مدل هایی که مبتنی بر یادگیری عمیق هستند مثل CNN، RNN و ترنسفورمهایی مثل BERT یا GPT، می توانند پیچیدگی های زبانی و زمینه ای را در ورودی های متنی درک کنند و با استفاده از روش هایی مانند تبدیل گفتار به متن و تحلیل احساسات متن، احساسات در دستگاه های IOT مانند اسپیکرهای هوشمند یا دستیارهای مجازی رو شناسایی کنند. این الگوریتم ها می توانند با تمرکز بر بخش های مهم محتوا، حتی نشانه های احساسی خیلی ظریف را هم تشخیص بدهند. [6] از آنجا که این ترنسفورمها از نظر محاسباتی خیلی سنگین هستند و باعث ایجاد مشکل در پردازش لحظه ای در محیط های IOT میشوند (به خصوص برای پردازش در خود دستگاه یا در لبه شبکه)، استفاده از مدل های تقطیر شده مانند DistilBERT گزینه بهتری است. مدل های تقطیر شده، مدل های بزرگ ترنسفورمر را به مدل های کوچکتر و سریع تر تبدیل

مدل	زمان آموزش	سرعت نتیجه گیری	میزان حافظه مصرفی	قابلیت استفاده بلادرنگ
Naive Bayes	کم	بسیار سریع	کم	مناسب برای وظایف ساده
SVM	متوسط	متوسط	متوسط	محدود به IoT با حجم بالا
LSTM	زیاد	کند	زیاد	با بهینه سازی قابل استفاده است
BERT	بسیار زیاد	متوسط تا کند	بسیار زیاد	فقط با رایانش لبه ای یا بهینه سازی
Logistic Regression	کم	سریع	کم	مناسب برای تحلیل احساسات پایه ای
RoBERTa	بسیار زیاد	متوسط تا کند	بسیار زیاد	در اپلیکیشن های بلادرنگ چالش برانگیز است
FastText	کم	سریع	متوسط	ایده آل برای محیط های بلادرنگ
TextBlob	بسیار کم	بسیار سریع	بسیار کم	مناسب تحلیل احساسات سبک و بلادرنگ

میکنند، بدون اینکه دقتشان کم بشود. با توجه به جدول ۱ نسخه‌ی تقطیرشده‌ی BERT بین سرعت بالا و دقت مناسب، تعادل برقرار می‌کند و برای تشخیص سریع احساسات در دستگاه‌هایی با توان پردازشی پایین بسیار مناسب است. [6]

جدول 1: مقایسه‌ی مدل‌های تحلیل احساسات از نظر کارایی و سازگاری با تحلیل بلادرنگ در محیط‌های IoT. [6]
مدل‌های ساده مانند Naive Bayes و Logistic Regression برای اپلیکیشن‌های IoT که سرعت و سادگی در آنها مهم‌تر از عمق مدل است، مناسب هستند. FastText دقت و سرعت بالایی دارد و با حافظه‌ی نسبتاً کم می‌تواند احساسات را به صورت لحظه‌ای تحلیل کند. در مقابل، BERT و RoBERTa دقت بسیار بالایی دارند، اما به دلیل نیاز شدید به توان پردازشی و حافظه، برای تحلیل لحظه‌ای در IoT چندان مناسب نیستند. مدل‌های LSTM هم برای داده‌های دنباله‌دار مناسب اند، اما در جریان‌های داده‌ای حجیم IoT سرعت تحلیل را پایین می‌آورند. [6]

۳. کاربردهای تحلیل احساسات

۳.۱. شهر هوشمند

تحلیل احساسات در شهرهای هوشمند با بهره‌گیری از داده‌های کاربران در شبکه‌های اجتماعی، ابزاری برای درک نگرش شهروندان است و مدیریت خدمات شهری را فراهم می‌کند. این تحلیل‌ها به ارتقای کیفیت زندگی شهری کمک می‌کنند. [3] تحلیل لحظه‌ای احساسات شهروندان از چند منبع و بهره‌گیری از الگوریتم‌های یادگیری عمیق، واکنش سریع به شرایط بحرانی را در شهرهای هوشمند فراهم می‌کند. تحلیل احساسات در شهرهای هوشمند کاربردهای متنوعی دارد مانند:

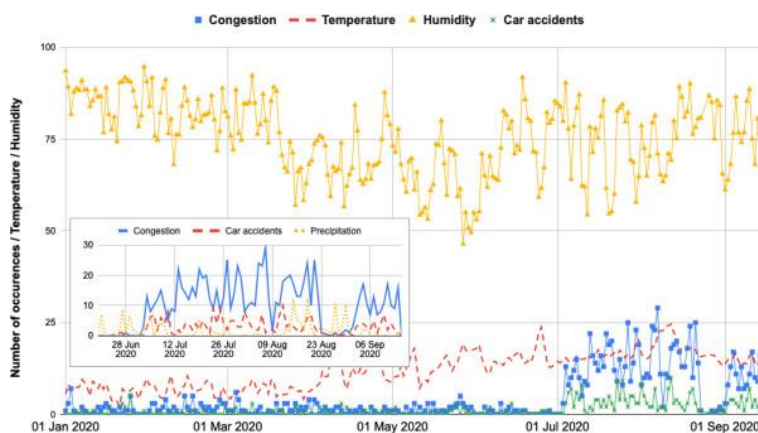
- مهم‌ترین کاربرد امکان دریافت بلادرنگ احساسات شهروندان در شبکه‌های اجتماعی است که به شناسایی ناآرامی‌های اجتماعی که در شکل ۴ به آنها اشاره شده کمک می‌کند. [3]



شکل ۴ - هفت دسته اصلی از رویدادهای موثر در شهر هوشمند [3]

- کاربرد بعدی درک میزان رضایت عمومی از خدمات شهری مانند حمل و نقل، کیفیت آب و هوا یا عملکرد نهادهای شهری است. سامانه‌هایی مانند "نبض شهری" با تحلیل احساسات کاربران در مناطق مختلف، وضعیت احساسی ساکنین هر منطقه را به تصویر می‌کشند. این ابزارها در تنظیم سیاست‌ها و خدمات شهری به مدیران کمک می‌کنند. [3]
- در بخش سلامت عمومی، از تحلیل احساسات در رسانه‌های اجتماعی برای ارزیابی واکنش عمومی نسبت به کمپین‌های واکسیناسیون و روندهای سلامت روان استفاده می‌شود. به عنوان مثال در طول همه‌گیری کووید-۱۹، پژوهشگران احساسات در توییتر را برای درک اضطراب عمومی و رعایت مقررات بهداشتی، تحلیل کردند. [9]

- در علوم سیاسی، از احساسات توییتر برای تحلیل حمایت عمومی از نامزدها و پیش‌بینی نتایج انتخابات می‌توان استفاده کرد. [9]
- کاربرد دیگر، ترکیب داده‌های محیطی با داده‌های احساسی شبکه‌های اجتماعی است که می‌تواند به شناسایی الگوهای جدید کمک کند. [3] در شکل ۵ تاثیر ترافیک، دما، رطوبت یا بارش بر میزان تصادفات به تصویر کشیده شده است.



شکل ۵ - نمونه ای از تاثیر عوامل مختلف بر میزان تصادفات [3]

تحلیل احساسات در بستر اینترنت اشیا یک حوزه نوظهور و کاربردی است که در آن تلاش می‌شود داده‌های خام سنسورها و دستگاه‌های متصل به اینترنت با اطلاعات مربوط به حالات احساسی کاربران ترکیب شوند تا درک عمیق‌تری از وضعیت کاربران و محیط اطراف حاصل شود. [6]

رویکردهای پردازش زبان طبیعی (NLP) برای پردازش داده‌های پیچیده و چندوجهی در محیط‌های اینترنت اشیا ضروری هستند، از جمله شاخص‌های احساسی مبتنی بر رفتار یا حسگرها. مدل‌های مدرن NLP می‌توانند ظرافت‌های زبانی و بافتی را در ورودی‌های متنی ثبت کنند، به‌ویژه معماری‌های یادگیری عمیق مانند شبکه‌های عصبی پیچشی (CNN)، شبکه‌های عصبی بازگشتی (RNN) و ترنسفورمرها مانند BERT یا GPT. این مدل‌ها می‌توانند متغیرهای مرتبط با احساسات را در قالب‌های غیرسنتی برای داده‌های IoT استخراج کنند. CNN الگوها را در داده‌های حسگرهای سری زمانی شناسایی می‌کند، در حالی که مدل‌های ترنسفورمر متن تعاملات گفتاری را تفسیر می‌کنند. به دلیل انعطاف‌پذیری‌شان، این الگوریتم‌ها می‌توانند نشانه‌های احساسی را هم در داده‌های متنی و هم غیرمتنی IoT تشخیص دهند. [6]

- اولین کاربرد در این زمینه شامل سیستم‌های هوشمندی است که توانایی تطبیق پاسخ‌ها بر اساس حالت‌های احساسی کاربران را دارند. برای مثال، دستگاه‌های هوشمند خانگی قادر خواهند بود با شناسایی احساس ناراحتی یا خستگی در کاربر، نور، دما یا صدای خانه را تنظیم کنند تا آرامش بیشتری فراهم شود. [6]
 - در حوزه سلامت نیز، داده‌های احساسی جمع‌آوری‌شده از طریق حسگرهای پوشیدنی می‌توانند برای پایش سلامت روان و جسم بیماران مورد استفاده قرار گیرند و به پزشکان در تشخیص و درمان مؤثرتر کمک کنند. [6]
 - یکی دیگر از کاربردهای کلیدی تحلیل احساسات در IoT، فراهم‌سازی امکان واکنش خودکار و متناسب با وضعیت احساسی در سامانه‌های بلادرنگ است. این واکنش‌ها ممکن است شامل ارسال هشدار، تغییر تنظیمات محیطی یا فعال‌سازی خدمات خاص باشند [6].
- در نهایت، ترکیب تحلیل احساسات با داده‌های دریافتی از سنسورها، مسیر آینده‌ی اینترنت اشیا را به سمت درکی عمیق‌تر از محیط و شکل‌گیری تعاملات هوشمند میان انسان و ماشین هدایت می‌کند. [6]

۳،۳. خدمات مشتری هوشمند

تحلیل داده‌های احساسی کاربران با استفاده از مدل‌های پیشرفته NLP و ترکیب داده‌های IoT، به شرکت‌ها اجازه می‌دهد خدمات شخصی‌سازی شده، پیش‌بینی نیازها و ارتقاء تجربه کاربری را به صورت بلادرنگ ارائه کنند که موجب افزایش رضایت و وفاداری مشتریان می‌شود.

این حوزه یکی از گسترده‌ترین و شناخته‌شده‌ترین حوزه‌های کاربرد تحلیل احساسات، در سیستم‌های خدمات مشتری مبتنی بر هوش مصنوعی است. تحلیل احساسات، تجربه مشتری را با افزایش شخصی‌سازی، کارایی و مقیاس‌پذیری متحول می‌کند. برخی از کاربردهای مستقیم شامل موارد زیر است:

- پیش‌بینی نارضایتی مشتری قبل از ترک پلتفرم یا شکایت رسمی

- پاسخ‌گویی احساسی توسط چت‌بات‌ها در زمان عصبانیت یا اضطراب مشتری
- تحلیل روند کلی احساسات در گزارش‌های خدماتی برای بهبود فرآیند
- پاسخ‌گویی به سوالات متداول و رفع مسائل رایج
- تعامل با مشتریان برای پیشنهاد محصولات [4]

تحقیقات بازاریابی برای ارزیابی رضایت مشتری و وفاداری به برند، با استخراج احساسات از بازخوردها و نقدهای آنلاین نتایج قابل توجهی ارائه می‌دهد. برای مثال، تغییرات در احساسات شناسایی شده در نقدهای محصول می‌تواند به عنوان شاخص‌های اولیه نارضایتی مصرف‌کننده عمل کند. [9]

در عصر تحول دیجیتال، تحلیل احساسات به عنوان یکی از ارکان اصلی در بهبود تجربه مشتری، اهمیت بسیاری پیدا کرده است. ادغام مدل‌های پردازش زبان طبیعی (NLP) با چارچوب‌های استاندارد خدمات مشتری نظیر مدل‌های SERVQUAL و Kano، مسیر نوینی برای توسعه خدمات مشتری هوشمند فراهم کرده است. الگوریتم‌های NLP می‌توانند با تحلیل داده‌های رفتاری و زبانی مشتری، این نیازها را دسته‌بندی کرده و به صورت پیش‌بینی‌پذیر، رضایت یا نارضایتی مشتری را شناسایی کنند. برای نمونه، چت‌باتی که با تحلیل احساسات و شناسایی واژگان کلیدی خشم یا نارضایتی، سریعاً گفتگو را به یک عامل انسانی منتقل می‌کند. [4]

۳،۴. اپلیکیشن‌های موبایل

در سالهای اخیر، رشد چشمگیر اینترنت بستر مناسبی برای افراد فراهم کرده تا نظرات و تجربیات خود را به اشتراک بگذارند. سهولت دسترسی و امکان ناشناس ماندن در این فضا، کاربران را به بیان دیدگاه‌های خود در پلتفرم‌هایی مانند فیس‌بوک، توییتر و آمازون تشویق می‌کند. هم‌زمان، شکل‌گیری پلتفرم‌های تجارت اجتماعی (Social E-commerce) نیز نقش مهمی در تعامل کاربران با محصولات و خدمات داشته‌اند. [2]

از جمله نمونه‌های برجسته این حوزه می‌توان به Red/XiaoHongShu در چین، Line Shopping در ژاپن، Amazon Live در آمریکا و TikTok Tokopedia Seller Center در اندونزی اشاره کرد. در این فضاها، کاربران به‌طور گسترده بازخورد خود را درباره تجربه‌های خرید، استفاده از خدمات و رضایت یا نارضایتی از محصولات ثبت می‌کنند. [2] تحلیل احساسات این امکان را فراهم کرده است که بازخوردهای متنی به‌صورت خودکار و ساختاریافته در قالب احساسات مثبت، منفی یا خنثی طبقه‌بندی شوند. برای انجام چنین تحلیلی از روش‌های مبتنی بر واژگان، به‌دلیل سادگی و قابلیت تفسیر بالا استفاده شده است. این روش‌ها از واژگان پیش تعریف‌شده، استفاده می‌کنند که کلمات را با امتیازهای احساسی مشخصی مرتبط می‌کنند. این روش‌ها در متون کوتاه مانند توییت‌ها و پست‌های فیس‌بوک، به‌ویژه هنگامی که تحلیل احساسات در سطح کلمه و بدون در نظر گرفتن بافت، کافی باشد، نتایج قابل اطمینانی ارائه داده‌اند. [9]

روش‌های پیشرفته‌تری مانند تحلیل معنایی نهفته (LSA) و مدل‌سازی موضوع با استفاده از تخصیص نهفته دیرایشلت (LDA) تلاش کرده‌اند تا ساختارهای موضوعی پنهان را استخراج کنند. با این حال، روش‌های سنتی غالباً در مواجهه با حساسیت به بافت و چندمعنایی بودن زبان غیررسمی شبکه‌های اجتماعی دچار مشکل می‌شوند. [9]

مدل‌های یادگیری عمیق به طور مؤثری روی مجموعه داده‌های شبکه‌های اجتماعی اعمال شده‌اند و عملکرد برتری در مدیریت متن‌های غیررسمی و بدون ساختار که در پلتفرم‌هایی مانند توییتر و ردیت رایج است ارائه می‌دهند. [9]

روش‌های پیشرفته پردازش زبان طبیعی و مدل‌های یادگیری عمیق مانند BERT و GPT، با توانایی درک معنای ضمنی و لحن گفتار کاربران، نه تنها طبقه‌بندی کلی احساس را ممکن می‌سازند، بلکه می‌توانند احساس کاربر نسبت به یک ویژگی خاص (مانند قیمت، کیفیت یا زمان ارسال) را در سطح عبارت یا جمله نیز تحلیل کنند. [4]

با توجه به نتایج حاصل، اپلیکیشن‌ها می‌توانند دیدگاه‌های کاربران را به صورت هدفمند در جهت بهبود عملکرد اپلیکیشن استفاده نمایند. تولیدکنندگان نیز با پایش مداوم نظرات کاربران نسبت به محصولات، قادر خواهند بود روند بازار را دقیق‌تر دنبال کرده و محصولات جدید را بر اساس نیازهای واقعی کاربران طراحی کنند. [2]

۴. چهارچوب پیشنهادی

در این مقاله هدف، تحلیل احساسات کاربران نسبت به دستگاه‌ها و خدمات مبتنی بر اینترنت اشیاء (IoT) است که تجربه خود را از طریق شبکه‌های اجتماعی به اشتراک می‌گذارند. برای این منظور، داده‌های متنی از پلتفرم‌هایی نظیر توییتر، آمازون، و بخش نظرات اپلیکیشن‌های هوشمند جمع‌آوری و پردازش شده‌اند. برای تحلیل موثر احساسات در بسترهای متنوعی مثل شهر هوشمند، داده‌های IOT، خدمات مشتری و اپلیکیشن‌های موبایل، یک چارچوب ترکیبی و ماژولار ارائه می‌شود که از قدرت ترکیبی یادگیری ماشین، NLP پیشرفته و معماری‌های بلادرنگ بهره می‌برد. داده‌ها شامل نظرات کاربران درباره دستگاه‌ها و خدمات متصل به IOT مانند ترموستات هوشمند، دوربین نظارتی، خانه‌های هوشمند، و اپلیکیشن‌های مرتبط و همچنین داده‌هایی از پلتفرم‌هایی مثل توییتر، آمازون، فیسبوک و TikTok Tokopedia Seller Center می‌باشند.

۴/۱. جمع‌آوری داده‌ها

داده‌ها با استفاده از ابزارها و منابع زیر استخراج شده‌اند:

- Twint: یک ابزار پایتونی برای جمع‌آوری توییت‌ها بدون نیاز به API رسمی توییتر [7]
- Tweet Retriever Coordinator: اسکریپت مدیریت‌شده برای جمع‌آوری توییت [7]

Twitter API : استفاده در سامانه‌هایی مانند TweetAlert و سامانه شناسایی رویدادهای تروریستی و سامانه تحلیل احساسات انتخابات برای شهر هوشمند [3]
Facebook API : جهت استخراج اطلاعات کاربران فیسبوک در سامانه ارزیابی احساسات [3]
Net-Twitter (Perl module) : در سامانه شناسایی رویدادهای تروریستی مبتنی بر تحلیل احساسات [7]
Google Play Store : جمع‌آوری نظرات کاربران اپلیکیشن TikTok Tokopedia Seller Center از طریق مکانیزم اسکرپینگ (Web Scraping) و پایتون [2]

۴/۲. پیش پردازش داده‌ها

در این مرحله از فرآیند، داده‌های خام و نامنظم پیش از تحلیل باید پاک‌سازی و آماده‌سازی شوند. پیش‌پردازش نقش کلیدی در تبدیل داده‌های اولیه بدون ساختار به داده‌های ساختاریافته و قابل استفاده برای مدل‌های یادگیری ماشین ایفا می‌کند. [2]

- پاک‌سازی متن : تکنیک‌هایی مانند شناسایی "not" و نگاشت شکلک‌ها به احساسات خاص می‌توانند دقت تحلیل احساسات را افزایش دهند. کاراکترهای خاص، URLها و تگ‌های HTML باید در این مرحله به درستی مدیریت شوند تا از تداخل با الگوریتم‌های تحلیل احساسات جلوگیری شود. تبدیل تمام متن به حروف کوچک می‌تواند به کاهش ابعاد فضای ویژگی‌ها کمک کند. تکنیک‌های اصلاح املا، مانند استفاده از واژه‌نامه‌های خاص زبان یا الگوریتم‌هایی مانند فاصله‌ی (Levenshtein distance) می‌توانند برای رفع مشکل اشتباهات املایی استفاده شوند. [5]
- حذف کلمات توقف : با حذف کلمات توقف (Stop words) مانند "the"، "and"، "is"، که در تحلیل احساسات معنای قابل توجهی ندارند، می‌توان نویز را کاهش داده و بهره‌وری محاسباتی را بهبود بخشید. [5]
- ریشه‌یابی : حذف پیشوندها یا پسوندها از کلمات برای تبدیل شان به ریشه اصلی مانند "Loving"، "Love" به "Love" در نرمال‌سازی متن و کاهش فضای ویژگی‌ها مفید هستند. [5]
- توکن‌سازی : شامل تقسیم متن به کلمات، عبارات یا جملات مجزا برای تسهیل تحلیل‌های بعدی است. تکنیک‌های رایج شامل توکن‌سازی مبتنی بر فاصله، توکن‌سازی مبتنی بر قواعد، و مدل‌های آماری می‌باشد. [5]
- پاک‌سازی نهایی : در این مرحله، داده‌های مخدوش شناسایی و اصلاح می‌شوند. چون داده‌های "کثیف" می‌توانند دقت تحلیل یا پیش‌بینی را تحت تأثیر قرار بدهند و این مرحله بسیار مهم است. [2]
- برچسب‌گذاری احساسات : در این مرحله، به هر داده برچسب احساسی مثبت، منفی یا خنثی داده می‌شود. این برچسب‌ها برای آموزش و ارزیابی مدل‌ها در مراحل بعدی مورد استفاده قرار می‌گیرند. [2]
- استخراج ویژگی‌ها : یک جمله یا سند به کلمات شکسته می‌شود؛ تا ماتریس ویژگی‌ها ساخته شود. مقدار هر خانه در ماتریس، تعداد تکرار آن کلمه در سند مربوطه را نشان می‌دهد. این ماتریس به مدل‌های یادگیری ماشین، ارائه شده و سپس عملکرد این مدل‌ها با استفاده از معیارهای ارزیابی مناسب سنجیده می‌شود. [1] این فرآیند به هر سند، بر

اساس محتوای آن، یک یا چند کلاس اختصاص می‌دهد. کلاس‌ها از ساختار طبقه‌بندی از پیش تعریف‌شده‌ای مانند سلسله‌مراتب دسته‌ها یا کلاس‌ها انتخاب می‌شوند. طبقه‌بندی متن از سناریوهای متنوعی پشتیبانی می‌کند، از جمله:

۱. طبقه‌بندی دودویی، مانند تحلیل ساده احساسات (مثبت یا منفی)
۲. طبقه‌بندی چندکلاسه، مانند انتخاب یک دسته از میان چند گزینه

اکثر الگوریتم‌های طبقه‌بندی مستقیماً متن خام را به‌عنوان ورودی دریافت نمی‌کنند؛ بلکه به بردارهای عددی نیاز دارند. برای این منظور، لازم است متن را به بردارهای دیجیتال قابل فهم برای الگوریتم‌ها تبدیل کنیم. [1]

TF-IDF: تکنیک TF-IDF برای تبدیل متن به نمایش‌های عددی به کار گرفته می‌شود و اهمیت واژه‌ها را با اختصاص وزن‌های بالاتر به واژگانی که در یک سند خاص پرتکرار هستند و در کل مجموعه داده نادر هستند، در نظر می‌گیرد. TF-IDF تأثیر واژگان پرتکرار را کاهش می‌دهد. [2] فرمول آن دو بخش دارد:

TF (Term Frequency): یعنی این کلمه چند بار در متن تکرار شده است.

IDF (Inverse Document Frequency): یعنی این کلمه چقدر بین بقیه متن‌ها نادر است.

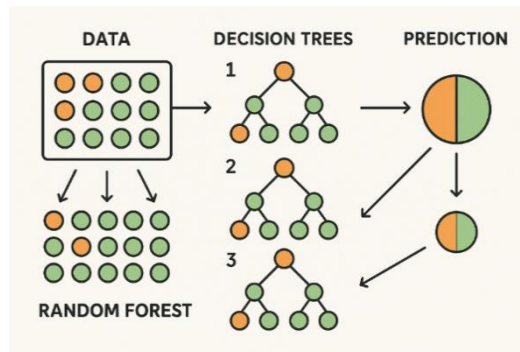
اگر یک کلمه در یک سند زیاد تکرار شده است ولی در بقیه اسناد خیلی کم آمده است، احتمالاً مهم است و TF-IDF بالایی دارد.

Term Presence: یکی از روش‌های ساده وزن‌دهی در اسناد متنی است که تنها بر وجود یا عدم وجود یک واژه در متن بدون توجه به تعداد تکرار آن تمرکز دارد. اگر واژه‌ای در فهرست ویژگی‌های مرجع قرار دارد، در متن مشاهده شود، مقدار (۱) و در صورتی که واژه مذکور در متن حضور نداشته باشد، مقدار (۰) به آن داده می‌شود. این روش در متونی با طول محدود (مانند توئیتهای یا نظرات کاربران) بسیار مؤثر است، به‌ویژه زمانی که حضور یک کلمه خاص به‌تنهایی نشانگر احساس مشخصی باشد. [2]

۴/۳. الگوریتم‌های طبقه‌بندی

در این مقاله از چندین الگوریتم برای دسته‌بندی داده‌های احساسات استفاده شده است:

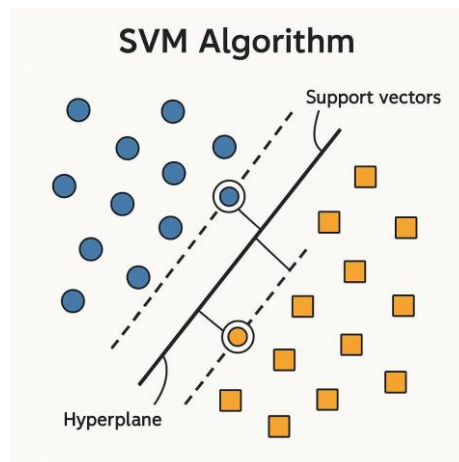
- **جنگل تصادفی (Random Forest)**: یک الگوریتم ترکیبی مبتنی بر مجموعه‌ای از درخت‌های تصمیم است که به خوبی از پس داده‌های نویزدار و نامتوازن برمی‌آید و به دلیل ساختار تصادفی و چندمدلی خود، مقاومت بالایی در برابر Overfitting دارد. [2] شکل ۶ مراحل کار الگوریتم جنگل تصادفی را نشان می‌دهد. مرحله اول ورود داده‌ها (شامل نمونه‌هایی با کلاس‌های مختلف نارنجی و سبز) میشود. در مرحله بعد چند درخت تصمیم، با استفاده از داده‌های تصادفی ساخته می‌شود و در نهایت هر درخت یک پیش‌بینی ارائه می‌دهد و با رأی‌گیری اکثریت، نتیجه نهایی مشخص می‌شود.



شکل ۶ - نحوه عملکرد الگوریتم Random Forest [8]

• ماشین بردار پشتیبان (SVM):

این الگوریتم یک یادگیری نظارت شده است که برای طبقه بندی و رگرسیون استفاده می شود. [1] SVM امکان استفاده از توابع کرنل مختلف را فراهم می کند که انعطاف پذیری لازم برای انطباق با انواع داده ها و حوزه های مسئله را فراهم می آورد. [5] شکل ۷ نحوه عملکرد SVM را نشان می دهد. دایره های آبی و مربع های نارنجی نمایانگر دو کلاس مختلف داده هستند که باید توسط الگوریتم از هم جدا شوند و خط وسط (Hyperplane) مرز تصمیم گیری است که داده ها را از هم جدا می کند. خط چین ها نشان دهنده حاشیه ها (Margins) هستند و داده هایی که دقیقاً روی این حاشیه ها قرار دارند به عنوان بردارهای پشتیبان شناخته می شوند و نقش مهمی در تعیین موقعیت Hyperplane دارند. هدف الگوریتم SVM این است که این hyperplane را طوری انتخاب کند که فاصله بین دو حاشیه (Margin) بیشینه شود، تا مدل قابلیت تعمیم دهی بهتری داشته باشد.



شکل ۷ - نحوه عملکرد الگوریتم SVM [8]

• رویکرد ترکیبی (Hybrid):

این روش برای حل مسئله ی طبقه بندی ، توانایی های دو الگوریتم Random Forest (RF) و Support Vector Machine (SVM) را به طور همزمان به کار می گیرد. [1] در این روش ورودی داده ها به هر دو مدل داده می شود و در نهایت خروجی دو مدل ترکیب شده و نتیجه به صورت یک رأی گیری یا میانگین گیری ارائه می گردد.

• نایو بیز (Naive Bayes) :

این الگوریتم در تحلیل نظرات کوتاه و ساده عملکرد خوبی دارد، اما در مواجهه با بازخوردهای طولانی یا پیچیده، چون فرض استقلال بین کلمات دارد، ممکن است دچار مشکل شود. این سیستم از یک رویکرد ترکیبی شامل سیستم های مبتنی بر قواعد و روش های خودکار استفاده می کند تا دقت محاسبه ی احساسات را افزایش دهد. [2]

۴/۴. ارزیابی مدل

معیار های ارزیابی عملکرد در تحلیل احساسات شامل Accuracy (دقت کلی)، Precision (دقت پیش بینی)، Recall (بازخوانی یا حساسیت) و F1-Measure (میانگین هارمونیک Precision و Recall) میشوند. این معیارها بر اساس مثبت یا منفی بودن بررسی ها توسط روش های بیان شده محاسبه و ارزیابی می شوند. [1] در این زمینه:

TP (True Positive): پیش بینی صحیح یک نتیجه مثبت

TN (True Negative): پیش بینی صحیح یک نتیجه منفی

FP (False Positive): پیش بینی اشتباه یک نتیجه مثبت (در حالی که واقعاً منفی بوده)

FN (False Negative): پیش بینی اشتباه یک نتیجه منفی (در حالی که واقعاً مثبت بوده)

این گزارش نسبت بین تعداد نمونه های مثبت که به درستی طبقه بندی شده اند و تعداد کل عناصری که باید به درستی طبقه بندی شوند را بیان میکند [1]

داده ی مربوط به نرخ های TP Rate (نرخ مثبت درست) و FP Rate (نرخ مثبت اشتباه) اجازه می دهد تا ماتریس آشفتگی (confusion matrix) برای یک کلاس خاص بازسازی شود. [1]

دقت پیش بینی (Precision) نسبت بین تعداد نمونه های مثبت درست (True Positive) و مجموع نمونه های مثبت درست و مثبت اشتباه (True Positive + False Positive) است. [1]

بازخوانی (Recall) درصد نمونه های مثبت واقعی است که به درستی شناسایی شده اند. [1]

دقت کلی (Accuracy) معیاری رایج برای ارزیابی عملکرد طبقه بندی است و نسبت نمونه های به درستی طبقه بندی شده به کل نمونه ها را نشان می دهد، در حالی که نرخ خطا (Error Rate) بر پایه تعداد نمونه های نادرست طبقه بندی شده تعریف می شود. [1]

امتیاز F1 (F1-Measure) این مقدار امکان جمع بندی عملکردهای طبقه بند (برای یک کلاس مشخص) را به صورت یک عدد واحد، با توجه به بازخوانی (Recall) و دقت (Precision) فراهم می کند. [1]

۵/۱. مقایسه دقت الگوریتم های یادگیری ماشین

در این قسمت عملکرد الگوریتم ها را در داده هایی که از پلتفرم های مختلف به دست آمده ، مقایسه میکنیم:
داده ها از شبکه ی اجتماعی توییتر استخراج شده و پس از بخش بندی و استفاده از تکنیک های حذف کلمات توقف (Stop Words) ، تبدیل به ۱۱۹۹۷ توییت (داده) شده است. برای آموزش مدل از ۸۰٪ توییت ها (داده ها) ۹۵۹۸ استفاده شده است که ۴۸۰۴ نمونه از آن ها مثبت و ۴۷۹۴ منفی بوده است . سپس با ۲۰٪ از توییت ها معادل ۲۳۹۹ توییت از عملکرد سیستم تست گرفته شده است که در نتیجه ۱۱۹۴ نمونه از داده ها مثبت و ۱۲۰۵ نمونه از داده ها منفی بوده است. [5]

مجموعه داده ها	داده های مثبت	داده های منفی	مجموع
آموزش داده شده	۴۸۰۴	۴۷۹۴	۹۵۹۸
تست شده	۱۱۹۴	۱۲۰۵	۲۳۹۹
مجموع	۵۹۹۸	۵۹۹۹	۱۱۹۹۷

بر اساس نتایج بدست آمده از پنج بار تکرار آزمایش، در جدول ۲ ، جدول ۳ و جدول ۴ مشخص است ، مدلی که با استفاده از الگوریتم ماشین بردار پشتیبان (SVM) ساخته شده، عملکرد بهتری نسبت به الگوریتم جنگل تصادفی (Random Forest) دارد. [5]

الگوریتم جنگل تصادفی نیز عملکرد قابل توجهی از خود نشان داده و دقت کلی معادل ۰.۷۸۵۶۴ به دست آورده است. [5]
در ادامه نتایج به دست آمده از آمازون را بیان میکنیم: برای این بخش ما از تکنیک هایی استفاده می کنیم که به طور خودکار داده ها را به احساسات مثبت یا منفی تقسیم می کنند. داده های آموزشی و آزمایشی این تحقیق شامل ۱۰۰۰ نمونه است؛ ۵۰۰ نمونه مثبت و ۵۰۰ نمونه منفی و داده خنثی نداریم. [1]

جدول ۲ - نتایج بدست آمده از داده های توییتر [5]

جدول ۳ - نتایج آزمون الگوریتم جنگل تصادفی و ماشین بردار پشتیبان را در پنج تکرار [5]

میانگین	آزمون ۵	آزمون ۴	آزمون ۳	آزمون ۲	آزمون ۱	الگوریتم ها
۰.۷۸۵۶۱	0.78112	0.79018	0.78596	0.78565	0.78518	Random Forest
۰.۸۰۲۲۲	0.80044	0.80680	0.80190	0.79971	0.80227	SVM

جدول ۴ - عملکرد الگوریتم SVM , Random Forest [5]

امتیاز F1	دقت پیش بینی (Precision)	یادآوری (Recall)	دقت کلی (Accuracy)	الگوریتم ها
0.78527	0.78737	0.78564	0.78564	Random Forest
0.80347	0.80654	0.80394	0.80394	SVM

در جدول ۵ مشاهده میکنیم که از ۱۰۰۰ بررسی، ۸۲۴ بررسی به درستی طبقه بندی شده و ۱۷۶ بررسی اشتباه طبقه بندی شده اند. [1]

جدول ۵- عملکرد الگوریتم SVM [1]

الگوریتم		مثبت	منفی	مجموع
SVM	مثبت	۴۰۹	۹۱	۵۰۰
	منفی	۸۵	۴۱۵	۵۰۰
	مجموع	۴۹۴	۵۰۶	۱۰۰۰

در جدول ۶ مشاهده میکنیم که از ۱۰۰۰ بررسی، ۸۱۰ بررسی به درستی طبقه بندی شده و ۱۹۰ بررسی اشتباه طبقه بندی شده اند. [1]

جدول ۶- عملکرد الگوریتم Random Forest [1]

الگوریتم		مثبت	منفی	مجموع
Random Forest	مثبت	۴۰۲	۹۸	۵۰۰
	منفی	۹۲	۴۰۸	۵۰۰
	مجموع	۴۹۴	۵۰۶	۱۰۰۰

در جدول ۷ از روش Hybrid یعنی ترکیب دو روش Random forest و SVM استفاده کرده ایم همینطور که می بینید از ۱۰۰۰ بررسی، ۸۳۴ بررسی به درستی طبقه بندی شده و ۱۶۶ بررسی اشتباه طبقه بندی شده اند.

جدول ۷- عملکرد الگوریتم RESVM [1]

الگوریتم		مثبت	منفی	مجموع
Hybrid)RFSVM (	مثبت	۴۱۶	۸۴	۵۰۰
	منفی	۸۲	۴۱۸	۵۰۰
	مجموع	۴۹۸	۵۰۲	۱۰۰۰

با توجه به جدول ۸ الگوریتم ترکیبی RFSVM، توانسته است ۸۳۴ نمونه را درست طبقه بندی کند و Precision، Recall و Accuracy آن برابر با ۸۳.۴٪ بوده است. که در مقایسه با سایر الگوریتمها نشان می دهد که RFSVM عملکرد بهتری دارد. رویکرد ترکیبی، مزایای هر دو الگوریتم Random Forest و SVM را ترکیب کرده و با استفاده از روش های یادگیری ماشین نظارت شده، دقت بالاتری ارائه می دهد و در مقایسه با سایر الگوریتمها، ثبات بیشتری دارد. [1]

جدول ۸ - تعداد نمونه طبقه بندی شده [1]

الگوریتم	نمونه های درست طبقه بندی شده (C.C.I)	نمونه های نادرست طبقه بندی شده (I.C.I)	دقت کلی (Accuracy)	دقت پیش بینی (Precision)	یادآوری (Recall)
SVM	824	176	82.4	82.4	82.4
Random Forest	810	190	81.0	81.0	81.0
RFSVM	834	166	83.4	83.4	83.4

در این قسمت نتایج به دست آمده از اپلیکیشن **TikTok Tokopedia Seller Center** را بیان میکنیم: قصد داریم احساسات کاربران نسبت به این اپلیکیشن را تحلیل کنیم برای این کار، از دیتاستی شامل 2000 نظر کاربر استفاده شده است و دو الگوریتم ماشین بردار پشتیبان (SVM) و نایو بیز (Naive Bayes) برای دسته بندی احساسات به سه گروه مثبت، منفی و خنثی به کار گرفته شده است. [2]

دیتاست مورد استفاده از طریق اسکرپ کردن نظرات اپلیکیشن در گوگل پلی با پایتون جمع آوری شده است. [2] بر اساس نتایج احساسات به دست آمده، مشخص شد که بیشترین احساس رایج، احساس منفی با تعداد کل ۱۱۷۱ مورد است، در حالی که احساس مثبت ۷۳۵ و احساس خنثی ۹۴ مورد بوده است.

در مرحله ی بعد، ارزیابی مدل طبقه بندی با استفاده از ماتریس آشفتگی (Confusion Matrix) انجام می شود تا کارایی و عملکرد دو الگوریتم بررسی گردد. در این مرحله با توجه به مقادیر جدول ۹ و جدول ۱۰ مشخص می شود که کدام روش وزن دهی (TF-IDF یا Term Presence) بهتر می تواند عملکرد مدل طبقه بندی را بهینه کند. [2]

جدول ۹- مقایسه نتایج ارزیابی مدل های طبقه بندی با استفاده از روش TF-IDF [2]

الگوریتم	احساسات	دقت کلی (Accuracy)	یادآوری (Recall)	دقت پیش بینی (Precision)	امتیاز F1
SVM	منفی	0.82	0.95	0.80	0.87
	خنثی		0	0	0
	مثبت		0.71	0.85	0.77
Naïve Bayes	منفی	0.69	0.99	0.66	0.79
	خنثی		0	0	0
	مثبت		0.32	0.94	0.48

جدول ۱۰- مقایسه نتایج ارزیابی مدل طبقه‌بندی با استفاده از روش Term Presence [2]

الگوریتم	احساسات	دقت کلی (Accuracy)	یادآوری (Recall)	دقت پیش‌بینی (Precision)	امتیاز F1
SVM	منفی	0.78	0.86	0.81	0.84
	خنثی		0	0	0
	مثبت		0.74	0.79	0.77
Naïve Bayes	منفی	0.75	0.84	0.77	0.80
	خنثی		0	0	0
	مثبت		0.71	0.72	0.73

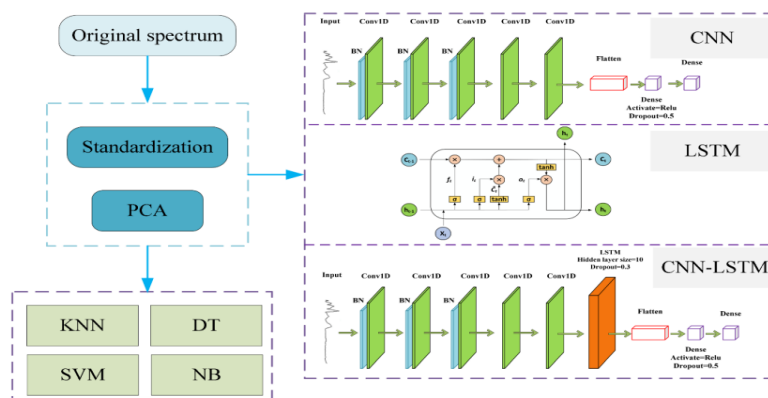
بر اساس نتایج ارزیابی برای مدل‌های طبقه‌بندی SVM و Naive Bayes، دقت مدل SVM چه در روش TF-IDF و چه در روش Term Presence بالاتر از مدل Naive Bayes بوده است. هر دو مدل در بخش احساس خنثی همچنان نتایج ضعیفی نشان داده‌اند. [2]

در نهایت، روش وزن‌دهی TF-IDF در بهینه‌سازی عملکرد مدل‌های طبقه‌بندی بسیار بهتر از روش Term Presence عمل کرده است. این به دلیل توانایی TF-IDF در شناسایی اطلاعات پیچیده‌تر و مرتبط‌تر در نظرات نسبت به Term Presence با نمایش باینری ساده آن می‌باشد. [2]

۵/۲. مقایسه الگوریتم‌های یادگیری عمیق

همانطور که در بخش ۲/۵ گفته شد مدل‌های شبکه عصبی بازگشتی برای پردازش داده‌های دنباله‌ای طراحی شده‌اند و پیشرفت بسیاری در این زمینه داشتند. مدل‌های LSTM برای غلبه بر مشکل کاهش گرادیان در RNN ها طراحی شدند و توانایی حفظ وابستگی‌های بلند مدت را به دست آوردند. [4]

RNN ها در وظایف احساسات توالی کوتاه با آموزش روی داده‌های کافی و استفاده از تعبیه‌های پیش‌آموزش‌دیده مانند Word2Vec یا GloVe نتایج قابل توجهی کسب کرده‌اند. مطالعاتی که از LSTM برای تشخیص احساس در سطح تویتر



استفاده کرده‌اند، بهبود دقت به‌ویژه در مدیریت طعنه، منفی‌سازی و بیان‌های هیجانی ظریف را گزارش کرده‌اند. [9] با توجه به شکل ۸ معماری‌های ترکیبی که CNN و LSTM را با هم به کار گرفته‌اند، عملکرد مدل را با درک هر دو الگوی محلی و ترتیبی در متن‌های حاوی احساس افزایش داده‌اند. [9] مدل‌های مبتنی بر LSTM در مجموعه داده‌هایی مانند SemEval, Yelp و پیکره‌های تویتر که وابستگی‌های زمانی و نحو پیچیده رایج است، به طور گسترده اعتبارسنجی شده‌اند. [9] GRUها در مدل‌سازی توالی‌های کوتاه و بلند نتایج رقابتی نشان داده‌اند و با هزینه محاسباتی کمتر برای کاربردهای بلنددرنگ تحلیل احساسات شبکه‌های اجتماعی مناسب هستند. [9] مطالعات مقایسه‌ای بین GRU و LSTM نشان داده‌اند که GRU در شرایط محدودیت داده آموزشی یا منابع اغلب عملکردی برابر دارد. [9]

شکل ۸- معماری‌های مدل برای تحلیل طیف با استفاده از شبکه عصبی کانولوشنی (CNN)، شبکه حافظه بلند-

کوتاه مدت (LSTM) و یادگیری ماشین کلاسیک [9]

مدل‌های شبکه عصبی بازگشتی به دلیل پردازش ترتیبی هنوز محدودیت‌هایی در گرفتن وابستگی‌های بلندمدت دارند. مدل‌های مبتنی بر ترنسفورمر برخلاف RNN ها، به جای پردازش ترتیبی، از مکانیزم Self-Attention برای درک همزمان وابستگی میان واژه‌های جمله بهره می‌برند که این توانایی ترنسفورمرها را در وظایفی که نیازمند فهم دقیق زبان است، بسیار موفق کرده است. [4]

مدل‌های ترنسفورمر باعث پیشرفت بیشتر در فهم و تولید زبان شده‌اند. به عنوان مثال همانطور که در بخش ۲/۵ گفته شد BERT متن را به صورت دوطرفه پردازش می‌کند که باعث میشود وابستگی‌های ظریف را درک کند و فهم پرسش‌ها را بهبود ببخشد. این درک عمیق متنی باعث افزایش دقت پاسخ‌های تولید شده توسط سیستم‌های هوش مصنوعی در خدمات مشتری می‌شود و آن‌ها را برای مدیریت پرسش‌های پیچیده و مبهم بهتر آماده می‌کند. [4]

همچنین رویکرد autoregressive در مدل‌های GPT باعث می‌شود در برنامه‌های مکالمه‌ای که حفظ جریان طبیعی گفت‌وگو اهمیت دارد، بسیار موثر باشد. توانایی GPT در تولید پاسخ‌های منسجم و حساس به زمینه، آن را به جزئی کلیدی در سیستم‌های هوش مصنوعی مکالمه‌ای تبدیل می‌کند که تعامل و رضایت مشتری را افزایش می‌دهد. [4]

این مدل‌های مبتنی بر ترنسفورمر به سیستم‌های هوش مصنوعی امکان می‌دهند تا حجم عظیمی از داده‌ها را پردازش کرده و از عهده وظایف زبان در مقیاس بزرگ برآیند. آنها می‌توانند مجموعه داده‌های آموزشی گسترده را به بخش‌های قابل مدیریت تقسیم کرده و به صورت موازی تحلیل کنند. این توانایی پردازش موازی، مقیاس‌پذیری را افزایش می‌دهد و به سیستم‌های هوش مصنوعی اجازه می‌دهد میلیون‌ها تعامل را همزمان بدون کاهش عملکرد، مدیریت کنند. همچنین قابلیت توزیع پردازش بین چند سرور یا نود باعث می‌شود زمان پاسخ حتی در دوره‌های ترافیک بالا، مانند تبلیغات فروش یا ساعات اوج کاری، پایین باقی بماند. [4]



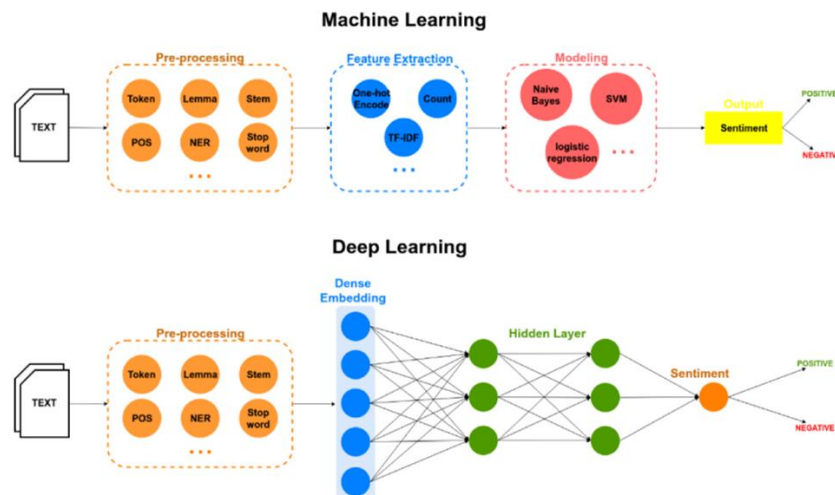
ارزیابی عملکرد مدل‌های مختلف تحلیل احساسات، از داده‌های متنی استخراج‌شده ی توئیتر، آمازون، توکوپدیا صورت گرفته است، که نشان می‌دهد الگوریتم SVM در مقایسه با Random Forest در داده‌های توئیتر در پنج بار تکرار و در مقایسه با Naive Bayes با استفاده از روش‌های وزن‌دهی TF-IDF و Term Presence در داده‌های متنی توکوپدیا، عملکرد بهتری نسبت به سایر الگوریتم‌ها داشته است. [5][1]

همچنین، نتایج کلی پژوهش نشان می‌دهد که استفاده از رویکردهای ترکیبی در مدل‌سازی و طبقه‌بندی احساسات، منجر به افزایش دقت، کاهش خطا و پایداری بیشتر در محیط‌های واقعی می‌شود. به‌ویژه، الگوریتم SVM در ترکیب با روش وزن‌دهی TF-IDF بهترین عملکرد را در طبقه‌بندی احساسات مثبت، منفی و خنثی بعد از مدل ترکیبی RFSVM ارائه داده است. [1]

الگوریتم‌های یادگیری عمیق، مدل‌های مبتنی بر ترنسفورمر از جمله GPT و BERT به دلیل درک همزمان وابستگی میان تمام واژه‌های جمله و توانایی حفظ وابستگی‌های بلند مدت نسبت به شبکه‌های عصبی بازگشتی عملکرد بهتری دارند و در تحلیل احساسات مشتری موفق‌تر عمل می‌کنند. [4]

الگوریتم‌های یادگیری ماشین قادر به پردازش و یادگیری حجم وسیعی از داده‌های حاشیه‌نویسی‌شده زبانی هستند. با این حال این مدل‌ها اغلب در مواجهه با پدیده‌های زبانی ظریف مانند طعنه، منفی‌سازی و عبارات اصطلاحی که نیازمند درک عمیق‌تر زمینه هستند، دچار مشکل می‌شوند؛ همچنین الگوریتم‌های یادگیری ماشین نظارت‌شده مانند ماشین بردار پشتیبان (SVM)، درخت تصمیم‌گیری و رگرسیون لجستیک برای طبقه‌بندی طعنه استفاده می‌شوند و بر ویژگی‌های دستی نظیر آن-گرام‌ها، برچسب‌های نقش کلمه (POS) و تضاد احساسی متکی هستند؛ با این حال، عملکرد این مدل‌ها به دلیل ناتوانی در ثبت وابستگی‌های بلندمدت و تفاوت‌های ظریف زبانی محدود است. معماری‌های یادگیری عمیق مانند CNN، LSTM، و مدل‌های مبتنی بر ترنسفورمر عملکرد بهتری دارند چرا که ویژگی‌های پنهان را از مجموعه داده‌های بزرگ استخراج می‌کنند. مدل‌های دوطرفه می‌توانند زمینه را از هر دو جهت دریافت کنند و در نتیجه در رفع ابهام عبارات طعنه‌آمیز بهتر عمل کنند. مدل‌های از پیش آموزش‌دیده و مبتنی بر ترنسفورمر عملکردی بهتر از روش‌های سنتی نشان داده‌اند. [9]

ویژگی‌های نحوی و معنایی در پرداختن به پدیده‌های زبانی مانند طعنه، کنایه، و منفی‌سازی که در گفتمان شبکه‌های اجتماعی رایج هستند، اهمیت ویژه‌ای دارند. برخلاف مدل‌های سنتی که به نمایش‌های پراکنده وابسته‌اند، معماری‌های یادگیری عمیق مانند RNN، LSTM، GRU می‌توانند این الگوهای معنایی و نحوی را از توالی‌ها استخراج کرده و دقت طبقه‌بندی را بهبود ببخشند. توانایی این مدل‌ها در یادگیری از بافت ترتیبی و درک وابستگی‌های عمیق ویژگی‌ها، آنها را برای کاربردهای غنی از احساس مانند شناسایی هیجان، تشخیص موضع و استخراج نظر مناسب ساخته است. [9] در شکل ۹ تفاوت یادگیری ماشین و یادگیری عمیق آورده شده است.



شکل ۹- تفاوت بین دو رویکرد طبقه‌بندی قطبیت احساسات؛ یادگیری ماشین و یادگیری عمیق [9]

مدل‌های مبتنی بر CNN در مجموعه داده‌های معیار از جمله SST، IMDb و نقدهای آمازون عملکرد بهتری نسبت به طبقه‌بندی سنتی یادگیری ماشین داشته‌اند. [9] چندین مطالعه گزارش کرده‌اند که مدل‌های یادگیری عمیق دقت طبقه‌بندی را بین ۸ تا ۱۵ درصد بیشتر از الگوریتم‌های سنتی بهبود داده‌اند. این مطالعات نشان داد که مدل‌های یادگیری عمیق و ترنسفورم‌ها توانسته‌اند در فهم ساختارهای پیچیده زبان و زمینه متنی بهتر عمل کنند.

۶. نتیجه‌گیری

در این پژوهش، با هدف طراحی یک چارچوب ترکیبی، جهت تحلیل احساسات کاربران در بستر اینترنت اشیاء اجتماعی، از داده‌های واقعی استخراج‌شده از شبکه‌های اجتماعی، اپلیکیشن‌های هوشمند و همچنین داده‌های رفتاری حاصل از دستگاه‌های متصل به اینترنت اشیاء استفاده گردیده است. [2] در این چهارچوب مدل‌های ذکر شده در مطالعات پیشین با یکدیگر ترکیب شده‌اند تا یک مدل جامع تر برای تحلیل احساسات کاربران در IOT اجتماعی ایجاد کنیم.

پس از جمع‌آوری داده‌ها به صورت دستی و پاکسازی داده‌ها، با استفاده از الگوریتم ترکیبی RFSVM و استفاده از مدل TF-IDF، نظرات مشتریان را تحلیل می‌کنیم. همچنین از مدل‌های مبتنی بر ترنسفورمر از جمله BERT و DISTILBERT، برای جمع‌آوری و تحلیل داده‌ها به صورت خودکار استفاده می‌کنیم. در این چهارچوب با توجه به مطالعات پیشین و در نظر گرفتن چالش‌هایی از قبیل زبان غیررسمی کاربران، ابهام معنایی در جملات و نیاز به تحلیل بلادرنگ در محیط اینترنت اشیاء، ترکیب مدل‌های مبتنی بر ترنسفورمر با الگوریتم‌های یادگیری ماشین می‌تواند راه‌حلی مؤثر برای درک بهتر رفتار کاربران و بهینه‌سازی سیستم‌های هوشمند باشد. ترکیب این مدل‌ها باعث می‌شود بتوانیم به جزئیات دقیق تری راجب نظر کاربر دست یابیم؛ به عنوان مثال دربابی که علت نارضایتی کاربر به دلیل عوامل محیطی (مانند گرمای هوا یا موقعیت جغرافیایی و یا مشکلات اجتماعی) بوده یا به دلیل نقص و مشکل در محصول و خدمات.



چارچوب طراحی شده در این تحقیق، قابلیت به کارگیری در حوزه های متعددی نظیر شهر هوشمند، خدمات مشتری، پایش سلامت روانی و تحلیل تجربیات کاربران را دارا است.
مهم ترین مزیت ها و یافته ها:

استفاده از مدل های ترکیبی پیشرفته مانند RFSVM (Random Forest Support Vector Machine) همراه با مدل های پیشرفته NLP مانند BERT و GPT، موجب افزایش قابل توجه دقت در تحلیل احساسات می شود. این ترکیب، توانایی بهتری در درک زمینه و پیچیدگی های زبانی فراهم می کند .
پردازش بلادرنگ داده های IoT با بهره گیری از الگوریتم های یادگیری عمیق و تکنیک های تحلیل چندمنبعی، امکان تصمیم گیری سریع و دقیق در سیستم های هوشمند را فراهم می آورد که به بهبود واکنش به شرایط بحرانی و بهینه سازی خدمات شهری کمک می کند. کاربرد هوش مصنوعی و پردازش زبان طبیعی در تحلیل احساسات کاربران در محیط های پیچیده و پویا، مانند شبکه های اجتماعی و شهرهای هوشمند، باعث می شود که سیستم ها بتوانند تغییرات سریع در رفتار و نیازهای کاربران را درک و پاسخگو باشند.

۶/۱. پیشنهادات برای پژوهش های آینده:

۱. استفاده از مدل های پیشرفته تری مانند یادگیری عمیق (Deep Learning) تا دقت طبقه بندی احساسات، به ویژه در بخش نظرات خنثی، بالاتر برود.
۲. استفاده از مدل های پیشرفته تر مانند BERT یا RoBERT برای درک بهتر زمینه معنایی در جملات پیچیده و اصطلاحات غیر رسمی.
۳. تحلیل چند وجهی احساسات با ترکیب داده های متنی، صوتی و تصویری منتشر شده توسط کاربران برای افزایش دقت.
۴. ساخت مجموعه داده های احساسی فارسی ویژه ی IOT برای کمک به آموزش دقیق تر مدل ها در بستر زبان فارسی.
۵. بکارگیری تحلیل بلادرنگ در لبه شبکه برای پیاده سازی مستقیم در محیط های عملیاتی IOT.
۶. بررسی تعامل احساسات کاربران با شرایط محیطی (دما، زمان، مکان) برای تحلیل رفتاری پیچیده تر.



منابع:

1. Al Amrani, Y., Lazaar, M., & El Kadiri, K. E. (2018). Random Forest and Support Vector Machine based Hybrid Approach to Sentiment Analysis. *Procedia Computer Science*, 127, 511–520. <https://doi.org/10.1016/j.procs.2018.01.150>
2. Dara, F. A., & Pratama, I. (2025). Sentiment Analysis of the TikTok Tokopedia Seller Center Application Using Support Vector Machine (SVM) and Naive Bayes Algorithms. *International Journal of Software Engineering and Computer Science*, 5(1), 177–189. <https://doi.org/10.35870/ijsecs.v5i1.3463>
3. Ben Ahmed, K., Radenski, A., Bouhorma, M., & Ben Ahmed, M. (2016). Sentiment Analysis for Smart Cities: State of the Art and Opportunities. In *Proceedings of the International Conference on Internet Computing and Internet of Things (ICOMP'16)* (pp. 55–58). CSREA Press. Retrieved from <https://worldcomp-proceedings.com/proc/p2016/ICM3084.pdf>
4. Chinna Raju, A. R. (2025). Advancing AI-Driven Customer Service with NLP: A Novel BERT-Based Model for Automated Responses. *Journal of Emerging Technologies and Innovative Research (JETIR)*, 12(1), 89–101. <https://www.jetir.org/view?paper=JETIR2501712>
5. Khan, T. A., Sadiq, R., Shahid, Z., Alam, M. M., & Su'ud, M. B. M. (2024). Sentiment analysis using Support Vector Machine and Random Forest. *Journal of Informatics and Web Engineering*, 3(1), 68–74. <https://doi.org/10.33093/jiwe.2024.3.1.5>
6. Manikyala, A. (2022). Sentiment analysis in IoT data streams: An NLP-based strategy for understanding customer responses. *Silicon Valley Tech Review*, 1(1), 35–47. <https://www.researchgate.net/publication/387427853>
7. Hodorog, A., Petri, I., & Rezgui, Y. (2022). Machine learning and Natural Language Processing of social media data for event detection in smart cities. *Sustainable Cities and Society*, 85, 104026. <https://doi.org/10.1016/j.scs.2022.104026>
8. OpenAI. (2025). Illustration of the Support Vector Machine (SVM) algorithm , combined function of SVM and Random Forest algorithms , how the Random Forest algorithm works [AI-generated image]. ChatGPT. <https://chat.openai.com/>
9. Alam, M. S., Mrida, M. S. H., & Rahman, M. A. (2025). Sentiment analysis in social media: How data science impacts public opinion knowledge integrates natural language processing (NLP) with artificial intelligence (AI). *American Journal of Scholarly Research and Innovation*, 4(1), 63–100. <https://doi.org/10.63125/r3sq6p80>