



مدل‌های زبانی بزرگ در کاوش داده‌های متنی

آزاده شاکری

دانشکده مهندسی برق و کامپیوتر - دانشگاه تهران

چکیده

در سال‌های اخیر، استفاده از شبکه‌های عصبی عمیق باعث پیشرفت چشم‌گیری در حوزه‌های مختلف از جمله بینایی ماشین، تشخیص گفتار، و کاربردهای مختلف پردازش زبان طبیعی شده است. از طرف دیگر مدل‌های زبانی بزرگ اخیراً پیشرفت شگفت‌انگیزی داشته‌اند و در مسائل مختلف ویژگی‌های قدرتمند و بی‌سابقه‌ای در استدلال، استنتاج و یادگیری مفاهیم پیچیده و سطح بالا از خود نشان داده‌اند. این مدل‌ها که با استفاده از حجم بسیار زیادی از داده‌ها، جزئیات زبان انسان را یاد می‌گیرند، به دلیل توانایی تولید متن منطقی و مناسب در زمینه‌های مختلف، به طور گسترده مورد استقبال قرار گرفته‌اند.

استفاده از مدل‌های زبانی بزرگ به فراتر از وظایف سنتی پردازش زبان طبیعی گسترش یافته است و شامل نمونه‌های گوناگونی از کاربردها می‌شود، از جمله تولید محتوا برای رسانه‌های اجتماعی، سیستم‌های توصیه‌گر شخصی سازی شده و کمک به تحقیقات علمی. علاوه بر این، مدل‌های زبانی بزرگ نقش مهمی در شکستن موانع زبانی ایفا کرده‌اند و ارتباطات چندزبانه را تسهیل کرده‌اند. کاربردهای متفاوت این حوزه در ادامه تشریح خواهند شد. موضوع این پژوهش، بهره‌گیری از قابلیت‌های مدل‌های زبانی بزرگ به منظور بهبود عملکرد در کاربردهای مختلف کاوش داده‌های متنی است.

از بین انواع مختلف داده‌ها، داده‌های متنی از اهمیت ویژه‌ای برخوردار هستند. اخیراً حجم داده‌های متنی بسیار زیاد شده است، و یک علت اصلی آن تکنولوژی‌هایی هستند که این امکان را فراهم می‌کنند که افراد به راحتی داده‌های متنی تولید کنند. به عنوان نمونه‌هایی از داده‌های متنی می‌توان از مقالات خبری، مقالات علمی، کتاب‌ها، کتابخانه‌های دیجیتال، ایمیل‌ها، صفحات وب، محتوای ایجاد شده در انجمن‌های پرسش و پاسخ، و محتوای متنی ایجاد شده در شبکه‌های اجتماعی نام برد. در این پژوهش انواع مختلف کاربردهای مرتبط با داده‌های متنی مورد توجه هستند.

کلیدواژه‌ها: مدیریت داده‌های متنی، متن‌کاوی، بازیابی اطلاعات، یادگیری عمیق، مدل‌های زبانی بزرگ

Large Language Models for Text Data Mining

Abstract

In recent years, the use of neural networks has made significant progress in various fields, including machine vision, speech recognition, and various applications of natural language processing. On the other hand, large language models have made amazing progress and have shown powerful and unprecedented features in reasoning, inference, and complex and high-level concepts in various problems. These models, which learn the details of human language by using vast amounts of data, have been placed in a general approach due to the production of logical and suitable text in different contexts.

The use of large language models has expanded beyond traditional natural language processing and includes a variety of examples of applications, including generating content for social media, personalized recommendation systems, and aiding scientific research. Moreover, language models have played an important role in the great linguistic break and created multilingual communication. The different applications of this field will be explained in the following. The subject of this research is to use the capabilities of large language models in order to improve the performance in various applications of text mining.

Among the different types of data, text is of particular importance. The volume of textual data has become enormous, and a major reason for that is the technologies that make it possible for people to easily generate textual data. Examples of textual data include news articles, scholarly articles, books, digital libraries, e-mails, web pages, created in question-and-answer forums, and text created in social networks. In this research, several applications related to textual data are of interest.

Keywords: Text data management, Text mining, Information retrieval, Deep learning, Large language models.

مقدمه و بیان مسئله

امروزه با افزایش روز افزون حجم داده‌های برخط رو به رو هستیم. دسترسی سریع و صحیح به منابع مهم و مورد علاقه، یکی از دغدغه‌های استفاده از منابع اطلاعاتی بسیار بزرگ است. در میان انواع منابع داده موجود، داده‌های متنی از اهمیت ویژه‌ای برخوردارند. متن طبیعی‌ترین روش رمزگذاری دانش بشری است. بعلاوه متن متداول‌ترین نوع اطلاعاتی است که مردم با آن روبرو می‌شوند. خیلی اطلاعاتی که افراد روزانه تولید و مصرف می‌کنند به صورت متن است. به عنوان نمونه می‌توان از ایمیل‌ها، پیام‌های متنی، مقاله‌ها، و پست‌های شبکه‌های اجتماعی نام برد. همچنین متن یک زبان نمایش عمومی است که برای توصیف رسانه‌های دیگر مانند فیلم یا تصاویر نیز استفاده می‌شود.

بر خلاف خیلی از سایر انواع داده‌ها، مانند داده‌هایی که توسط حسگرها جمع آوری می‌شوند، داده‌های متنی توسط انسان‌ها و برای استفاده انسان‌ها ایجاد می‌شوند. این داده‌ها به طور کلی دارای محتوای معنایی غنی هستند و اغلب حاوی دانش، اطلاعات، نظرات و ترجیحات با ارزش افراد هستند، که فرصت خوبی برای کشف انواع دانش مفید برای بسیاری از کاربردها فراهم می‌کنند. به عنوان مثال، افراد برای بدست آوردن نظرات درباره موضوعات جالب و بهینه‌سازی تصمیم‌گیری‌های مختلف مانند انتخاب یا خرید یک محصول، از داده‌های متنی نظیر بررسی محصولات، بحث‌های انجمن‌ها و متن رسانه‌های اجتماعی استفاده می‌کنند. به دلیل حجم گسترده اطلاعات، مردم به ابزارهای نرم‌افزاری هوشمند برای کمک به کشف دانش مربوطه برای بهینه‌سازی تصمیمات یا کمک به آنها در انجام کارها با کارایی بیشتر، نیاز دارند. فناوری‌های پشتیبانی از متن‌کاوی در سال‌های اخیر پیشرفت چشمگیری داشته‌اند و ابزارهای متن‌کاوی امروزه به طور گسترده در بسیاری از حوزه‌های کاربردی مورد استفاده قرار می‌گیرند.

متن‌کاوی تمام فعالیت‌هایی که به نوعی به دنبال کسب دانش از متن هستند را شامل می‌گردد. تحلیل داده‌های متنی توسط روش‌های یادگیری ماشین، بازیابی هوشمند اطلاعات، پردازش زبان طبیعی یا روش‌های مرتبط دیگر همگی در حوزه متن‌کاوی قرار می‌گیرند. متن‌کاوی به دنبال استخراج اطلاعات مفید از داده‌های متنی غیرساخت‌یافته از طریق تشخیص و نمایش الگوها است یا به عبارت دیگر روشی برای استخراج دانش از متون است. متن‌کاوی کشف اطلاعات جدید و از پیش ناشناخته، به وسیله استخراج خودکار اطلاعات از منابع مختلف نوشتاری است.

در طول دهه گذشته، بهبودهای چشم‌گیری در حوزه‌های بینایی ماشین، تشخیص گفتار، ترجمه ماشینی و متن‌کاوی هم در پژوهش‌ها و هم به صورتی کاربردی ایجاد شده است [۱]. این پیشرفت‌های گسترده به واسطه معرفی مدل‌های شبکه عصبی جدید با چندین لایه مخفی که با عنوان شبکه‌های عمیق مطرح شده‌اند، صورت گرفته است. همچنین کاربردهای نوینی از جمله عامل‌های سخن‌گو و عامل‌های بازی نیز ظهور پیدا کرده‌اند [۲، ۳]. امروزه پژوهش‌های وسیعی بر روی شبکه‌های عصبی عمیق در متن‌کاوی و بازیابی اطلاعات صورت می‌گیرد.

در یک تعریف کلی، یادگیری عمیق، همان یادگیری ماشین است، به طوری که در سطوح مختلف نمایش یا انتزاع یادگیری را برای ماشین انجام می‌دهد. با این کار، ماشین درک بهتری از واقعیت وجودی داده‌ها پیدا کرده و می‌تواند الگوهای مختلف را شناسایی کند. یادگیری عمیق در ابتدا توسط Hinton در سال ۲۰۰۶ به عنوان یک شبکه عصبی عمیق که کار یادگیری ماشین انجام می‌دهد، معرفی شد [۴]. شبکه‌های عصبی الهام‌گرفته از مغز انسان و شامل نورون‌ها و لایه‌های مختلفی است.

معماری عمیق این شبکه‌ها شامل سطح‌های مختلفی از عملیات غیرخطی است. این شبکه‌ها قادر به آموزش به دو شکل با نظارت و بدون نظارت هستند. از اوایل سال‌های ۲۰۰۰ شبکه‌های عصبی برای پیاده‌سازی مدل‌های زبانی مورد استفاده بوده‌اند [۵].

در طول سالیان، مدل‌های زبانی به عنوان ابزار قدرتمندی برای وظایف پردازش زبان طبیعی به وجود آمده‌اند. یکی از انواع مختلف این مدل‌ها، مدل‌های زبانی بزرگ هستند که با استفاده از حجم بسیار زیادی از داده‌ها، جزئیات زبان انسان را یاد می‌گیرند. این مدل‌ها به دلیل توانایی تولید متن منطقی و مناسب در زمینه‌های مختلف، به طور گسترده مورد استقبال قرار گرفته‌اند.

استفاده از مدل‌های زبانی بزرگ در سال‌های اخیر به فراتر از وظایف سنتی پردازش زبان طبیعی گسترش یافته است و شامل نمونه‌های گوناگونی از کاربردها می‌شود، از جمله تولید محتوا برای رسانه‌های اجتماعی، سیستم‌های توصیه‌گر شخصی سازی شده و کمک به تحقیقات علمی. انعطاف‌پذیری و قابلیت تطبیق این مدل‌ها آنها را در صنایع مختلف بسیار مهم کرده است و پیشرفت‌های در زمینه چت‌بات‌ها، دستیارهای مجازی و تولید محتوای خودکار حاصل شده است. علاوه بر این، مدل‌های زبانی بزرگ نقش مهمی در شکستن موانع زبانی ایفا کرده‌اند و ارتباطات چندزبانه را تسهیل کرده‌اند. کاربردهای متفاوت این حوزه در ادامه تشریح خواهند شد.

کاربردهایی از کاوش و بازیابی داده‌های متنی و پژوهش‌های پیشین

◀ بازیابی پرسش بین‌زبانی در انجمن‌های پرسش و پاسخ

امروزه اطلاعات موجود در انجمن‌های پرسش و پاسخ در کنار اطلاعات موجود در وب به کمک جستجوگران آمده و بسیاری از افراد نیازهای اطلاعاتی خود را در این انجمن‌ها مرتفع می‌کنند. در بسیاری از مواقع پرسش‌های کاربران تکراری و پاسخ‌های آنها در انجمن موجود است و کافیت که پرسش مشابه پرسش کاربر بازیابی شود و پاسخ آن مورد استفاده قرار گیرد. در شرایطی که پرسش ورودی به یک زبان و پرسش‌هایی که بازیابی باید از میان آنها صورت گیرد به زبان دیگری باشند مسئله تبدیل به مسئله بازیابی پرسش بین‌زبانی خواهد شد.

یک رویکرد برای حل مسئله بازیابی پرسش بین‌زبانی استفاده از بازنمایی کلمات مبتنی بر بافتار است. در روش‌هایی مانند BERT بازنمایی از پیش آموزش دیده شده برای کلمات هم در محیط تک‌زبانه و هم در محیط دوزبانه وجود دارد، اما برای استفاده در یک تسک خاص نیاز به روش‌هایی برای میزان‌سازی کردن آن به منظور استخراج ویژگی‌های خاص تسک مورد بررسی است. به کارگیری روش BERT در بسیاری از تسک‌ها از جمله پرسش و پاسخ، بازیابی اطلاعات، کاربرد در محیط دو زبانه و بازیابی اطلاعات بین زبانی که مرتبط با موضوع رساله می‌باشند انجام شده است. یک راه معمول برای استفاده از مدل BERT در تسک‌های رتبه‌بندی، میزان‌سازی آن با کمک طبقه‌بندی است با این هدف که آیا زوج ورودی (زوج پرس و جو-سند، زوج پرسش-پاسخ، زوج پرسش-پرسش یا ...) مرتبط هستند یا خیر. در ادامه از امتیاز بدست آمده برای رتبه‌بندی استفاده می‌شود. اما نتایج آزمایش‌های انجام شده در نشان می‌دهد که الگوریتم‌های رتبه‌بندی مبتنی بر جفت که توانایی تشخیص برتری یک سند بر دیگری را دارند یا الگوریتم‌های رتبه‌بندی مبتنی بر لیست که توانایی

بهینه‌سازی لیست را دارند نیز می‌توانند در تسک‌های رتبه‌بندی عملکرد مناسبی داشته باشند. در این پژوهش انواع روش‌های یادگیری رتبه‌بندی و استفاده از روش‌های بازنمایی کلمات مبتنی بر بافتار به منظور بازیابی پرسش در انجمن‌های پرسش و پاسخ با هم ترکیب شده‌اند. از آنجا که یکی از نقاط مورد تاکید در پیشنهاد پژوهشی استفاده از فراداده موجود در انجمن است، روش فوق روی بخش پاسخ پرسش‌ها نیز به طور مستقل اعمال می‌شود و در نهایت نتایج حاصل از فیلدهای پرسش و پاسخ با هم ترکیب می‌شوند [۱۱-۱۳].

◀ استفاده از نظرات کاربران برای استنباط دانش ذهنی در سیستم‌های توصیه‌گر مکالمه‌ای

دسته‌ای از سوالاتی که کاربران قبل از خرید یک محصول یا رزرو یک خدمت با آن‌ها مواجه می‌شوند، سوالاتی هستند که در پایگاه‌های داده شامل داده‌های واقعی مثل اندازه، قیمت، امتیاز و مشخصات محصول مورد نظر وجود ندارند و مربوط به ویژگی‌های ذهنی هستند. به عنوان مثال، در مورد کیفیت سرویس‌دهی یک هتل یا رستوران، شرایط لغو رزرو، میزان امنیت یک خودرو یا مناسب خانواده بودن اتمسفر یک رستوران، مواردی هستند که پاسخ آن‌ها را باید در نظرات، تجربیات و پرسش‌های پرتکرار کاربران در مورد محصول یا خدمات مورد نظر جستجو کرد. کاری که افراد برای دستیابی به پاسخ مورد نظرشان می‌کنند این است که قبل از خرید آن محصول، به سراغ بررسی‌ها و نظرات دیگر کاربران رفته و از تجربیات آن‌ها برای این منظور استفاده می‌کنند.

این دانش‌ها که مبنای آن تجربیات، نظرات و احساسات افراد است را به دانش ذهنی تعبیر کرده و قصد داریم در این پژوهش به استنباط این دانش‌ها از سخنان کاربران در سیستم‌های توصیه‌گر مکالمه‌ای پرداخته تا این سیستم‌ها بتوانند پاسخ مناسبی برای این دسته از سوالات کاربران داشته باشند.

روش کار و معماری سیستم توصیه‌گر مکالمه‌ای مبتنی بر دانش ذهنی^۱ به صورتی که در ادامه معرفی می‌شود، خواهد بود.

مرحله‌ی اول:

ابتدا با توجه به محتوای مکالمه و آخرین سخن کاربر، تشخیص داده می‌شود که آیا برای پاسخ به این گفته‌ی کاربر نیاز به استنباط دانش ذهنی وجود دارد و یا می‌توانیم از سیستم‌های توصیه‌گر مکالمه‌ای مرسوم برای تولید پاسخ مناسب استفاده کنیم. چالشی که در این مرحله هنوز هم در مقالات وجود دارد، قابلیت تعمیم این تشخیص برای وضعیت‌ها و دامنه‌های دیده‌نشده است که قصد داریم از قابلیت تعمیم و دانش نهفته در مدل‌های زبانی بزرگ برای این مسئله‌ی طبقه‌بندی دوکلاسه و طبقه‌بندی گفته‌های کاربران استفاده کنیم.

مرحله‌ی دوم:

پس از تشخیص نیاز به استنباط دانش ذهنی، دانش‌های مرتبط با محتوای مکالمه و گفته‌ی کاربر، از منابع دانش ذهنی و واقعی^۲ استخراج شده و نوعی رتبه‌بندی برای استخراج این برش‌های دانش صورت می‌گیرد. در این مرحله نیز چالش

¹ Subjective

² Factual

تعمیم‌پذیری سیستم مطرح خواهد بود. روشی که در مقالات برای این مسئله ارائه شده بود داده‌افزایی با استفاده از مدل‌های زبانی بزرگ بوده و مشکلی که وجود دارد، لزوم استناد برش‌های دانش بر واقعیت است که این روش از داده‌افزایی تا الان نتوانسته این تضمین را فراهم کند. مخصوصاً این که در این مسئله نظرات کاربران مطرح است، نیاز به داده‌های واقعی خواهد بود. به همین دلیل، قصد داریم از داده‌ها و نظرات کاربران درباره‌ی موجودیت‌ها (در این مسئله، هتل‌ها و رستوران‌ها) که به صورت عمومی در سطح وب قرار دارند برای غنی‌سازی منبع دانش ذهنی استفاده کنیم.

مرحله‌ی سوم:

در نهایت، گزیده‌ای از برش‌های دانش استخراج‌شده یا چکیده‌ای از آنها را به عنوان پاسخ سیستم تولید می‌کنیم. برای ارزیابی پاسخ تولید شده، می‌توانیم از روش‌های خودکار و یا روش‌هایی که از مدل‌های Pre-trained برای ارزیابی پاسخ‌ها استفاده می‌کنند و وابسته به بافتار هم هستند استفاده کنیم. با توجه به ماهیت تعاملی سیستم‌های توصیه‌گر مکالمه‌ای، قصد داریم از روش‌های شبیه‌سازی کاربر هم برای ارزیابی پاسخ‌های تولیدشده توسط سیستم استفاده کنیم. با دانش حال حاضر ما نمونه‌ای از کاربرد این روش‌ها برای مسئله‌ی استنباط دانش ذهنی در این سیستم‌ها مشاهده نشده و برای ارزیابی پاسخ‌ها از معیارهایی که گفته شد و یا ارزیابی انسانی استفاده می‌شود.

◀ سامانه‌های مکالمه‌ای جویای اطلاعات

طراحی و ساخت سامانه‌هایی که قابلیت گفتگو با انسان را دارند از چالش‌های اصلی علم هوش مصنوعی است و چنین سامانه‌هایی تکامل بعدی رابط‌های کامپیوتر-انسانی هستند. دسترسی به مجموعه داده‌های مکالمه‌ای بزرگ و پیشرفت‌های اخیر در یادگیری عمیق پیاده‌سازی سامانه‌های مکالمه‌ای مبتنی بر داده را امکان‌پذیر کرده‌اند. سامانه مکالمه‌ای یا سامانه گفتگو، به سامانه‌های هوشمندی گفته می‌شود که با استفاده از واسطه‌های کاربری مختلف مانند صوت یا متن و در قالب مکالماتی به زبان طبیعی با انسان تعامل داشته باشد، به گونه‌ای که سامانه در هر زمان با توجه به تاریخچه تعاملات کاربر با سامانه تا آن لحظه و درخواست فعلی کاربر، مناسب‌ترین پاسخ را به کاربر ارائه دهد. اگرچه با وجود پیشرفت‌های روز افزون در حوزه یادگیری عمیق و بهره‌گیری از آن در سامانه‌های پردازش زبان طبیعی طراحی این سامانه‌ها نیز رشد چشمگیری داشته است، با این حال همچنان توانایی این سامانه‌ها در ایجاد یک مکالمه پیوسته و روان با انسان با محدودیت‌های زیادی مواجه است. بنابراین برقراری یک سامانه مکالمه‌ای خودکار بین انسان و رایانه یکی از مسائل مهم و قابل توجه در علوم کامپیوتر به خصوص هوش مصنوعی است که روز به روز بر اهمیت آن افزوده می‌شود.

سامانه‌های مکالمه‌ای را با توجه به هدفی که برای آن طراحی شده‌اند، می‌توان به سه دسته تقسیم کرد: (۱) سامانه‌های پرسش و پاسخ مکالمه‌ای، (۲) سامانه‌های مبتنی بر وظیفه و (۳) سامانه‌های غیر وظیفه‌گرا یا همان چت‌بات‌ها. در سامانه‌های پرسش و پاسخ کاربران می‌توانند درخواست‌های خود را به صورت سؤالاتی به زبان طبیعی از یک پایگاه دانش بزرگ یا مجموعه‌ای از اسناد مشخص مطرح کنند و سامانه در هر مرحله با توجه به مجموعه سوال‌ها و جواب‌های قبلی و همچنین منبع دانش مورد نظر به سوال کاربر پاسخ دهد. حالت ساده‌تر این نوع از سامانه‌ها، همان ماشین‌های درک مطلب هستند که با توجه به متنی که در اختیار سامانه قرار گرفته است قادرند به سؤالات مرتبط با آن متن پاسخ دهند. در این

نوع از سامانه‌ها هر سوال کاربر به طور مجزا و مستقل از پرسش و پاسخ‌های قبلی، توسط سامانه پردازش شده و پاسخ مرتبط از روی متن استخراج می‌شود. دسته دوم، سامانه‌های مبتنی بر وظیفه هستند که در آن سامانه به کاربر کمک می‌کند تا بتواند وظیفه مشخصی را انجام دهد. برای مثال پیدا کردن یک محصول خاص و خرید اینترنتی آن، رزرو هتل یا رستوران، خرید بلیط هواپیما و سفارش غذا نمونه‌هایی از وظایفی است که این نوع از سامانه‌ها کاربران را در انجام آن راهنمایی می‌کنند. سامانه‌های غیر وظیفه گرا یا همان چت‌بات‌ها، با هدف پشبرد مکالمه با کاربر با ارائه پاسخ‌هایی منطقی، جملات خبری، جملات احساسی، طرح سوالات و یا درخواست اطلاعات بیشتر از کاربر طراحی می‌شوند. در این سامانه‌ها معمولا مکالمات بر روی دامنه باز بوده و موضوعات مطرح شده در آن‌ها محدوده مشخصی ندارد، مانند توییتر و Weibo. در این میان نوع خاصی از مکالمات هستند که در آن کاربران به دنبال اطلاعات خاصی هستند و تمام مکالمات انجام شده در آن در راستای رفع نیاز اطلاعاتی کاربر است. در این نوع از سامانه‌ها که مبتنی بر اطلاعات هستند دستیار مکالمه‌ای در صورتی که سوال کاربر نامفهوم باشد با طرح سؤالاتی سعی می‌کند اطلاعات بیشتری از کاربر جمع‌آوری نماید. سپس با دریافت پاسخ‌های کاربر و تحلیل آن‌ها در نهایت اطلاعات مورد نیاز کاربر را در اختیارش قرار می‌دهد. به این نوع از سامانه‌ها، سامانه‌های مکالمه‌ای جستجوی اطلاعات گفته می‌شود. در این پژوهش تمرکز بر روی طراحی این نوع از سامانه‌های مکالمه‌ای است.

یکی از ویژگی‌های بارز مکالمه حافظه‌دار بودن آن است، به این معنی که می‌توان به گفته‌های قبلی ارجاع کرد [۱۴]. یکی از بزرگ‌ترین کاستی‌های سامانه‌های موجود در صنعت عدم توانایی آن‌ها در نگهداری بافتار مکالمه است [۱۵]. یک سامانه‌ی کارآمد باید بتواند از گفته‌های قبلی کاربر در رفع نیاز اطلاعاتی کاربر استفاده کند و کاربر باید این قابلیت را داشته باشد که بتواند به گفته‌های قبلی خود یا سامانه ارجاع دهد. برای بهره‌گیری از اطلاعات موجود در این بافتار باید روشی برای مدل‌سازی آن در نظر گرفته شود.

مدل‌سازی بافتار در زمینه‌های مختلفی از بازیابی اطلاعات و پردازش زبان طبیعی کاربرد دارند. یک نمونه در چت‌بات‌های جویش اطلاعات مانند چت‌بات پشتیبانی مشتری است. برای آموزش این چت‌بات‌ها از آرشیو بزرگ مکالمه‌های قبلی بین دو انسان استفاده می‌شود و چت‌بات می‌تواند برای پاسخ کاربر از این آرشیو یک پاسخ مناسب انتخاب کند یا یک جمله‌ی جدید تولید کند [۱۶]. یک نمونه‌ی دیگر از کاربردهای مدل‌سازی بافتار در سامانه‌های سوال و جواب چند نوبتی هستند. در این سامانه‌ها به کاربر قابلیت سوال پرسیدن راجع به یک متن یا یک گراف دانش داده می‌شود و سامانه باید به کاربر پاسخی به زبان طبیعی برگرداند [۱۷].

سامانه‌های مکالمه‌ای را با توجه به روشی که برای ارائه پاسخ استفاده می‌کنند می‌توان به طور کلی به دو دسته مبتنی بر تولید پاسخ و مبتنی بر بازیابی پاسخ تقسیم نمود. در مدل مبتنی بر تولید، سامانه با تحلیل مکالمات قبلی کاربر با رایانه سعی بر تولید یک پاسخ مناسب با درخواست فعلی کاربر را دارد. این در حالی است که در سامانه‌های مبتنی بر بازیابی، بهترین پاسخ از میان مجموعه‌ای از پاسخ‌های کاندید موجود انتخاب می‌شود. اخیرا مطالعات زیادی در این حوزه صورت گرفته است [۱۸-۲۱] که در آن‌ها تاثیر روش‌های یادگیری عمیق در بهبود کارایی این سامانه‌ها بررسی شده است.

◀ تشخیص پیام‌های توهین‌آمیز در فضای بین‌زبانی با مدل‌سازی گرافی

تولید محتوای توهین‌آمیز^۳ که در سال‌های اخیر در ارتباطات برخط افزایش یافته است، جرم محسوب می‌شود. عوامل مختلفی در گسترش جملات تنفرآمیز^۴ و توهین‌آمیز وجود دارد. از یک طرف در اینترنت و به طور خاص شبکه‌های مجازی، به علت ماهیت مخفی که این محیط‌ها ایجاد می‌کنند، مردم رفتار خشونت‌آمیز بیشتری از خود نشان می‌دهند. از طرف دیگر، مردم تمایل بیشتری به بیان عقایدشان به صورت برخط دارند، در نتیجه انتشار محتوای توهین‌آمیز تسهیل می‌شود. از آنجا که این نوع تعاملات متعصبانه در بستر شبکه‌های اجتماعی می‌تواند برای جامعه مضر باشد، سکوهاى شبکه اجتماعى می‌توانند از ابزارهای تشخیص و پیش‌گیری در این مورد بهره‌برند.

محتوای تنفرآمیز پدیده پیچیده‌ای است که به طور ذاتی مربوط به روابط درون‌گروهی می‌شود و با ظرافت‌های زبان ارتباط دارد. این پدیده از دید شبکه‌های اجتماعی مختلف، تعاریف تا حدی نزدیک و البته متفاوتی دارد. در یک تعریف جامع، محتوایی تنفرآمیز قلمداد می‌شود که بیان تهاجمی یا هدف تحقیر داشته باشد به طوری که خشونت و نفرت را ضد گروه‌ها بر اساس ویژگی‌های خاص نظیر ظاهر فیزیکی، مذهب، ملیت، گرایش جنسی، هویت جنسی و سایر موارد که ممکن است با زبان متفاوتی حتی به شکل ظریفی با طنز بیان شود، گسترش می‌دهد. تصمیم‌گیری در مورد اینکه آیا بخشی از متن حاوی محتوای تنفرآمیز است، حتی برای انسان‌ها نیز ساده نیست. در نتیجه نیازمند ارائه راه‌کارهایی برای بهبود دقت روش‌های تشخیص این گونه محتواها هستیم [۲۲].

در کارهای پژوهشی پیشین، مفاهیم تا حدی مرتبط با این موضوع تحت عنوان آزار و اذیت سایبری^۵، زبان سوء استفاده^۶، تبعیض^۷، دشنام^۸ و افراطی‌گرایی^۹ آورده شده است [۲۳، ۲۴]. هر کدام از این مفاهیم از جهتی به پیام‌های توهین‌آمیز و تنفرآمیز شبیه و از جهت خاصی نیز متفاوتند. بررسی هر کدام از مفاهیم مرتبط و چالش‌های مربوط به دسته‌بندی آن‌ها، دقت مدل‌ها را افزایش می‌دهد. همچنین تشخیص موارد فوق در فضایی که کاربران به چندین زبان پیام ارسال می‌کنند نیز از چالش‌های دیگر موجود در این حوزه است.

در اکثر روش‌های ارائه شده برای تشخیص خودکار سخنان نفرت‌انگیز و جملات توهین‌آمیز به بررسی مساله در یک زبان و با استفاده از ویژگی‌های زبانی پرداخته شده است. عمده روش‌های ارائه شده از یک روش یادگیری ماشین برای استخراج و دسته‌بندی محتوای توهین‌آمیز استفاده کرده‌اند. در حالی که تحقیقات در این زمینه به سرعت در حال افزایش است، یکی از چالش‌های مهم آن است که بیشتر روش‌ها برای چند زبان غنی از منابع پیشنهاد شده‌اند. از طرفی در چند پژوهش سعی شده است تا با استفاده از روش‌های یادگیری انتقالی، دانش مدل‌های آموزش شده در زبان‌های غنی از منابع را به زبان‌های کم‌منابع تعمیم دهند. حال آنکه روش‌های یادگیری انتقالی لزوماً در همه موارد باعث بهبود دقت طبقه‌بندی بر

³ Offensive

⁴ Hate speech

⁵ Cyberbullying

⁶ Abusive language

⁷ Discrimination

⁸ Profanity

⁹ Extremism

روی داده زبان کم‌منبع نبوده‌اند. در این پیشنهاد پژوهشی، روش‌های مبتنی بر مدل‌سازی داده با گراف چندلایه و ناهمگون و سپس یادگیری بازنمایی از روی ساختار گراف ارائه شده است. در یک رویکرد داده‌های هر زبان به طور مستقل توسط گراف مدل‌سازی شده و در یک چارچوب مبتنی بر بازسازی یال‌های درون‌لایه‌ای و میان‌لایه‌ای، سعی در بهبود یادگیری بازنمایی‌ها شده است. در هر لایه گراف، یال‌هایی از جنس ارتباط بین کلمه و پیام، کلمه و کلمه و هشتگ و پیام در نظر گرفته شده است. ارتباطات بین‌لایه‌ای از طریق اتصال میان کلمات توهین‌آمیز معادل، در زبان‌های مختلف ایجاد می‌شوند. جهت غنی‌تر کردن ارتباطات بین‌لایه‌ای، استفاده از یک پایگاه دانش، به عنوان پیشنهاد مطرح شده است. در رویکردی دیگر، پیشنهاد یادگیری یال‌های میان‌لایه‌ای ارائه شده است تا به نوعی ساختار گراف روی هر مجموعه داده را کامل‌تر کند. به منظور ارزیابی روش‌های پیشنهادی، از چندین مجموعه داده در زبان‌های غنی از منابع و کم‌منابع استفاده شده است. نتایج اولیه آزمایش‌ها نشان می‌دهد که روش‌های پیشنهادی، دقت طبقه‌بندی پیام‌های توهین‌آمیز در زبان کم‌منبع را افزایش داده‌اند.

◀ مدل‌سازی کاربران برای پیش‌بینی واکنش در شبکه‌های اجتماعی با کمک مدل‌های زبانی

بزرگ

شبکه‌های اجتماعی برخط، بستری برای برقراری ارتباط، تعامل و تولید محتوای کاربران هستند. به دنبال رشد سریع فناوری و دسترسی آسان به اینترنت، حضور و فعالیت کاربران در شبکه‌های اجتماعی رشد چشم‌گیری داشته است. کاربران با اهداف مختلفی چون تعامل و برقراری ارتباط با سایر کاربران، تولید محتوا، بیان احساسات و نظرات، پیگیری اخبار و کاوش اطلاعات جدید در موضوعات مختلف از شبکه‌های اجتماعی بهره می‌برند. به دست آوردن مدلی مناسب برای فهم و پیش‌بینی رفتار کاربران، اهمیت بسیاری در تحلیل‌ها و وظایف مختلف پژوهشی شبکه‌های اجتماعی دارد و حجم بالایی فعالیت و داده‌های تولید شده، چالش‌ها و فرصت‌های فراوانی را در این زمینه به دنبال داشته است [۲۵].

به طور کلی پژوهش‌های حوزه رفتاری شبکه‌های اجتماعی را می‌توان به دو دسته تقسیم کرد؛ رفتار انفرادی کاربر و رفتار دسته جمعی کاربران به عنوان یک شبکه که در پی تعامل و ارتباطات کاربران پدید می‌آید. به دست آوردن مدلی کارا از کاربر که بتواند با توجه به مواردی مثل شرایط شبکه، سابقه کاربر، ویژگی‌ها و ترجیحات شخصی او، رفتار و واکنش‌هایش را پیش‌بینی کند از اهمیت بالایی در مسائل متعدد برخوردار است. برخی از این مسائل عبارتند از پیش‌بینی الگوی پخش اخبار در شبکه، شناسایی و جلوگیری از پخش اخبار جعلی، بهبود سامانه‌های پیشنهاد دهنده، تبلیغات بهینه و بازاریابی [۲۶].

مدل‌های زبانی بزرگ^۱ در سال‌های اخیر پیشرفت شگفت‌انگیزی داشته‌اند و در مسائل مختلفی مانند سامانه‌های پیشنهاد دهنده، ویژگی‌های قدرتمند و بی‌سابقه‌ای در استدلال، استنتاج و یادگیری مفاهیم پیچیده و سطح بالا از خود نشان داده‌اند. لازم به ذکر است که بسیاری از این کاربردها، نیازی به آموزش مدل با داده‌های ویژه آن وظیفه ندارند. توانایی‌های این

¹ Recommender Systems	0
¹ Large Language Models	1

مدل‌ها در استدلال علی، منطقی و قیاسی فرصت‌های جدیدی را برای یادگیری مفاهیم پیچیده و نهان در ساختارهای گرافیکی و همچنین سازوکارهای شناختی انسان ایجاد کرده است [۵۱-۵۲].

موضوع این پژوهش، بهره‌گیری از قابلیت‌های ذکر شده از مدل‌های زبانی بزرگ به منظور بهبود عملکرد در پیش‌بینی واکنش کاربر در شبکه‌های اجتماعی و پاسخ به تمام یا بخشی از چالش‌های ذکر شده است.

◀ خبره‌یابی

مسئله‌ی خبره‌یابی تا به حال در زمینه‌های مختلفی همچون پیشنهاد دکتر به بیماران [۲۷]، یافتن استاد راهنمای مناسب برای دانشجویان [۲۸] و مسیریابی سوال [۲۹] مورد استفاده قرار گرفته است. اما یکی از مهم‌ترین کاربردهای آن انتخاب افراد مناسب از بین مجموعه‌ای از داوطلبان برای یک موقعیت شغلی [۳۰] داده شده می‌باشد. در اینجا به طور کلی هدف این است که با در اختیار داشتن لیستی از حوزه‌های مورد نیاز یک موقعیت شغلی، یک رتبه‌بندی از داوطلبان بر اساس میزان دانششان در حوزه‌های ذکر شده تولید شود. تا به حال تحقیقات زیادی در این زمینه انجام شده است [۳۱-۳۲]، اما بسیاری از آنها قادر به برآورده کردن نیازهای اصلی شرکت‌ها نیستند و بدون توجه به نیازهای اصلی آنها اقدام به بازیابی افراد خبره می‌نمایند.

امروزه بسیاری از شرکت‌ها به دنبال استخدام افراد خبره‌ای هستند که بتوانند بر اساس متدلوژی‌های توسعه‌ی نرم‌افزار همچون متدلوژی چابک در قالب یک تیم بر روی یک پروژه داده شده فعالیت نمایند. در یک تیم ایده‌آل که مبتنی بر متدلوژی چابک است، تعدادی افراد خبره‌ی تی-شکل وجود دارد، به طوریکه هر یک در یک حوزه‌ی مورد نیاز تیم خبره و در سایر حوزه‌های مورد نیاز دانش عمومی دارد [۳۳]. اخیراً تعدادی تحقیق در این زمینه انجام شده است. در این تحقیقات، هدف اصلی، یافتن یک رتبه‌بندی از افراد خبره‌ی تی-شکل است که در یک حوزه‌ی خاص خبره باشند [۳۶-۳۴]. در اینجا به منظور ساده‌سازی، توجهی به دانش عمومی کاربران نشده و مسئله‌ی بازیابی تیم‌های چابک نادیده گرفته شده است. رستمی و نشاطی [۳۷] برای اولین بار به مسئله‌ی تشکیل یک تیم چابک با بازیابی افراد خبره‌ی تی-شکل پرداختند. آنها با استفاده از مدل‌های احتمالاتی مولد که مبتنی بر تعداد اسناد داوطلبان در حوزه‌های مورد نیاز تیم بودند اقدام به بازیابی تیم‌های چابک نمودند. با این حال، آنها هیچ توجه‌ای به محتوای اسناد داوطلبان نداشتند و دانش داوطلبان را بدون توجه به محتوای اسنادشان تخمین زدند.

رویکرد ما برای حل این مسئله به این صورت است که ابتدا با استفاده از یک روش پیشنهادی، بهترین داوطلبان ممکن را برای هر کدام از جایگاه‌های تیم انتخاب می‌کنیم. سپس با استفاده از یک مدل برنامه‌ریزی خطی عدد صحیح، بهترین تیم چابک ممکن را از بین داوطلبان انتخاب شده تشکیل می‌دهیم. برای آنکه بتوان بهترین داوطلبان ممکن را برای هر جایگاه از تیم پیش‌بینی کرد، مدل پیشنهادی ما بایستی بتواند به طور همزمان از دانش حاصل از شکل خبرگی و سطح دانش تخصصی و عمومی هر داوطلب در حوزه‌های مهارتی مورد نیاز جایگاه، برای تخمین مرتبط بودن او به آن جایگاه استفاده کند. ما تا به حال با دو رویکرد به طراحی چنین مدلی پرداخته‌ایم:

رویکرد اول-مدل رمزگذار چندگانه: در اینجا از یک مدل رمزگذار استفاده می‌شود که با یک رویکرد یادگیری با نظارت به تخمین احتمال مرتبط بودن هر داوطلب می‌پردازد. این مدل در واقع از ترکیب بازنمایی‌های برداری حاصل از تعدادی رمزگذار، برای تخمین احتمال ذکرشده استفاده می‌کند. در اینجا، از یک رمزگذار شکل خبرگی برای تولید بازنمایی مربوط به شکل خبرگی داوطلب و از تعدادی رمزگذار سطح دانش برای تولید بازنمایی‌هایی که نشان دهنده سطح دانش تخصصی/عمومی داوطلب در حوزه‌های مهارتی مورد نیاز باشد، استفاده می‌شود. برای طراحی رمزگذار شکل خبرگی، ما از یک شبکه عصبی بازگشتی که از فرآیند توزیع تاپیک‌های اسناد داوطلب در گذر زمان بهره می‌برد، استفاده کرده‌ایم تا بتوانیم پروفایل داوطلب را به‌صورت یک بازنمایی برداری تولید کنیم. همچنین برای طراحی رمزگذارهای سطح دانش، بازنمایی برداری پروفایل داوطلب را با بازنمایی مربوط به هر حوزه‌ی مهارتی مورد نیاز تیم مقایسه کرده‌ایم تا دانش داوطلب در هر حوزه را به‌صورت یک بازنمایی برداری تخمین بزنیم.

رویکرد دوم- مدل احتمالاتی: در اینجا با یک رویکرد یادگیری بدون نظارت به تخمین احتمال مرتبط بودن هر داوطلب می‌پردازیم. این مدل در واقع با الهام‌گیری از مشخصه‌های اصلی یک تیم چابک، احتمال مرتبط بودن داوطلب به یک جایگاه از تیم را بسته به برقراری همزمان سه شرط پوشش، ارتباط و بهینگی برای حضور در آن جایگاه می‌داند. در اینجا، برای تخمین احتمال برقراری شروط پوشش و ارتباط، از یک رمزگذار سطح دانش و برای تخمین احتمال برقراری شرط بهینگی، از یک رمزگذار شکلی خبرگی استفاده می‌شود. ما در این رویکرد، برای طراحی رمزگذار سطح دانش، از یک رویکرد مبتنی بر توجه استفاده کرده‌ایم که با بررسی میزان توجه یک حوزه‌ی مهارتی مورد نیاز تیم به پروفایل داوطلب، یک بازنمایی شخصی‌سازی شده از حوزه‌ی مهارتی را ایجاد می‌کند. سپس با مقایسه‌ی این بازنمایی با پروفایل داوطلب، ما به تخمین سطح دانش داوطلب در آن حوزه می‌پردازیم. همچنین برای طراحی رمزگذار شکل خبرگی، ما ویژگی‌های شکل خبرگی داوطلب را با تحلیل مجموعه بازنمایی‌های برداری مربوط به دانش داوطلب در حوزه‌های مهارتی یاد می‌گیریم، سپس با داشتن این ویژگی‌ها، توزیع شکل خبرگی داوطلب را تخمین می‌زنیم.

◀ ارائه رویکرد چند اظهاره جهت بررسی لزوم پیشنهاد سوال شفاف‌سازی و انتخاب سوال

بهبود در سیستم‌های مکالمه

موتورهای جستجو اهمیت بی‌ش از پیش در زندگی روزمره افراد پیدا کرده است و افراد به صورت روزمره از این موتورها استفاده می‌کنند. موتورهای جستجو وبسایت‌های مورد نظر کاربر را با توجه به پرس‌وجوی مطرح شده از سوی کاربر از مجموعه وبسایت‌های موجود در وب بازیابی می‌کنند و به افراد ارائه می‌دهند. گاهی پرس‌وجوی کاربر کوتاه یا مبهم است و درک پرس‌وجوی کاربر به صورت کامل اتفاق نمی‌افتد و موتورهای جستجو با روش‌های متفاوتی از جمله تنوع بخشی به نتایج [۴۵] یا دسته‌بندی نتایج [۴۶] مانند آنچه در موتور جستجوی Clusty رخ می‌دهد- تلاش می‌کردند تا این مشکل را کم‌رنگ کنند. روشی جدید که به منظور حل مشکل مذکور مورد توجه قرار گرفته است، استفاده از پرسش شفاف‌ساز [۴۷] است. در این روش موتور جستجو با ارائه پرسشی در پاسخ به پرس‌وجوی مبهم و مشخص کردن تعدادی پاسخ کاندید برای این پرسش، تلاش می‌کند تا با دریافت اطلاعاتی اضافی از کاربر و گسترش پرس‌وجوی کاربر با استفاده از اطلاعات دریافت شده، نتایج بهتری را ارائه کند. یکی از دلایل توجه روز افزون به این شکل برخورد با پرسش‌های مبهم، رواج استفاده

از سیستم‌های جستجوگرهای مکالمه‌ای وب است. در این ابزارها، تنها یک نتیجه به کاربر بازگردانده می‌شود و روش‌های مرسوم مورد استفاده در موتورهای جستجو سنتی قابل استفاده نیستند. همچنین در صورت عدم شفاف‌سازی پرس‌وجوی مبهم و ارائه یک نتیجه، امکان دقیق نبودن و کاربردی نبودن آن برای کاربر وجود دارد.

این روش مانند هر روش دیگری، چالش‌هایی در پیاده‌سازی و اجرا دارد. یکی از این چالش‌ها تشخیص پرس‌وجوهای مبهم [۴۸] است تا در زمان مناسب پرس شفاف‌ساز ارائه شود. اگر چه در طی تحقیقات مشخص شده است که کاربران از برقراری تعامل با سیستم لذت می‌برند، اما در صورتی که پرسش شفاف‌ساز در پاسخ به یک پرس‌وجوی واضح ارائه شود برای کاربر دلسرد کننده است و در صورت تداوم پرسیدن پرسش‌های شفاف‌ساز در یک جستجوگر مکالمه‌ای وب، منجر به ترک مکالمه از سوی کاربر می‌شود و می‌تواند به اندازه نرسیدن این پرسش و ارائه پاسخ اشتباه به کاربر، برای سیستم ضرر به همراه داشته باشد. چالشی دیگر بازیابی پرسش مناسب [۴۹] برای پرس‌وجوی مطرح شده است. انتخاب پرسش مناسب از بین مجموعه پرسش‌های موجود در سیستم، اهمیت بالایی دارد چرا که در صورت انتخاب پرسش نامناسب مکالمه از مسیر خود منحرف می‌شود و به نتایج صحیح نمی‌انجامد.

هدف ما ایجاد بهبود در بازیابی پرسش شفاف‌ساز از بین مجموعه‌ای بسیار بزرگ از این پرسش‌ها است. در تلاش هستیم تا با تشخیص حوزه پرس‌وجوی مطرح شده، و استفاده از آن در کنار پرس‌وجو، برای تشخیص ابهام و همچنین به منظور بازیابی پرسش شفاف‌ساز مناسب، به نتایجی بهتر از آنچه در مطالعات قبلی [۴۹] به آن دست یافته‌اند برسیم. در این مسیر از مجموعه داده‌ای با نام MIMICS [۵۰] که از داده‌های پرسش شفاف‌ساز موجود در موتور جستجو بینگ جمع آوری شده است، استفاده کنیم.

➤ تشخیص سخنان توهین‌آمیز و تنفرآمیز در شبکه‌های اجتماع با در نظر گرفتن زمینه‌های

موثر بر کاربر

ظهور سکوه‌های رسانه‌های اجتماعی برخط تحولی عمیق در نحوه برقراری ارتباط، اشتراک‌گذاری اطلاعات و تعامل افراد با یکدیگر ایجاد کرده است. استفاده از این فضاها به بخشی از زندگی روزمره ما تبدیل شده‌اند و بستری را برای بیان نظرات، ارتباط با دیگران و تقویت جوامع فراهم می‌کنند. استفاده از سکوه‌های شبکه اجتماعی در دهه گذشته به طور تصاعدی رشد کرده‌اند و شبکه‌های اجتماعی سریع‌ترین ابزار ارتباطی هستند زیرا پیام‌ها تقریباً بلافاصله ارسال و دریافت می‌شوند. با این حال، در کنار مزایای آن‌ها، نگرانی فزاینده‌ای در مورد شیوع و تأثیر سخنان توهین‌آمیز و تنفرآمیز در این سکوها وجود دارد.

افزایش استفاده از سکوه‌های رسانه‌های اجتماعی برخط، افزایش مشابهی را در شیوع سخنان توهین‌آمیز و تنفرآمیز ایجاد کرده است و چالش‌های مهمی را برای افراد و جامعه به همراه داشته است لذا تشخیص و رسیدگی سریع به این موارد ضروری است. تشخیص خودکار این نوع سخنان چالش‌های زیادی دارد، ماهیت پویای زبان و تفاوت‌های ظریف زبانی، شناسایی خودکار سخنان توهین‌آمیز و تنفرآمیز را پیچیده‌تر می‌کند. این چالش‌ها نیازمند رویکردهای نوآورانه‌ای هستند که می‌توانند ماهیت متنوع این نوع از سخنان را در زمینه‌های مختلف توضیح دهند.

بیشتر پژوهش‌های انجام شده در حوزه تشخیص سخنان توهین‌آمیز و تنفرآمیز صرفاً به محتوای متنی اکتفا کرده‌اند و در این میان، کمتر توجهی به مزایای ساختار تعاملات بین کاربران و زمینه‌های موثر بر آنان شده است. همچنین بیشتر پژوهش‌های انجام شده بر روی زبان انگلیسی تمرکز کرده‌اند و کمتر پژوهشی به زبان فارسی پرداخته است. در این پژوهش، ما رویکرد جدیدی را پیشنهاد می‌کنیم که بر اهمیت تعاملات کاربر و عوامل زمینه‌ای موثر بر کاربران، در تشخیص سخنان توهین‌آمیز و تنفرآمیز تأکید می‌کند. با تجزیه و تحلیل الگوهای تعاملات کاربر و تأثیر شبکه‌های اجتماعی، می‌توانیم بینش عمیق‌تری در مورد پویایی انتشار سخنان توهین‌آمیز و تنفرآمیز به دست آوریم. علاوه بر این، در نظر گرفتن عوامل زمینه‌ای امکان درک دقیق‌تری از این نوع از سخنان را فراهم می‌کند و دقت مدل‌های تشخیص را افزایش می‌دهد. در این پژوهش مجموعه داده‌ای به زبان فارسی ایجاد خواهیم کرد که در آن ساختار تعاملات بین کاربران نیز در نظر گرفته شده است و مدل پیشنهادی با این مجموعه داده آموزش داده خواهد شد.

در حال حاضر، به منظور صحت‌سنجی و ارزیابی اولیه مدل پیشنهاد شده، اقدام به جمع‌آوری زیرمجموعه‌ای از داده‌های شبکه اجتماعی پاساژ کرده‌ایم و مجموعه داده‌ای برچسب‌زنی شده که شامل متن و ساختار تعاملات بین کاربران است را آماده کرده‌ایم. بر روی این مجموعه داده مدل‌های پایه آزمایش و اجرا گردید و عملکرد آن‌ها ارزیابی شد. براساس ایده مطرح شده، اقدام به بررسی عوامل موثر در بهبود عملکرد تشخیص این نوع سخنان کرده. طبق بررسی‌ها یکی از مواردی که می‌تواند به بهبود عملکرد مدل‌های تشخیص سخنان تنفرآمیز و توهین‌آمیز کمک کند در نظر گرفتن تعاملات بین کاربران است، برای بررسی این ادعا آزمایشی بر روی مجموعه داده ایجاد شده انجام گردید و درستی این ادعا مورد بررسی قرار گرفت. از دیگر مواردی که مورد بررسی قرار گرفت بررسی اثر سوابق فردی کاربران در تشخیص سخنان تنفرآمیز بود، با انجام آزمایشی بر روی مجموعه داده نمونه نشان داده شد که در صورتی که سوابق کاربر در نظر گرفته شود عملکرد مدل بهبود پیدا می‌کند.

جمع بندی

در این پژوهش، تمرکز بر روی استفاده از مدل‌های زبانی بزرگ در چند کاربرد مختلف کاوش داده‌های متنی و بازبینی اطلاعات است. به طور خاص در این پژوهش، سامانه‌های مکالمه‌ای جویای اطلاعات، تشخیص محتوای توهین‌آمیز در فضای چندزبانه، تشخیص سخنان توهین‌آمیز و تنفرآمیز در شبکه‌های اجتماعی با در نظر گرفتن زمینه‌های موثر بر کاربر، مدل‌سازی کاربران برای پیش‌بینی واکنش در شبکه‌های اجتماعی با کمک مدل‌های زبانی بزرگ، رویکرد چند اظهاره جهت بررسی لزوم پیشنهاد سوال شفاف‌سازی و خبره‌یابی مورد بررسی قرار خواهد گرفت.

مراجع

1. LeCun, Y., Y. Bengio, and G. Hinton, *Deep learning*. Nature, 2015. 521: p. 436.
2. Vinyals, O. and Q. Le, *A neural conversational model*. arXiv preprint arXiv:1506.05869, 2015.

3. Silver, D., Huang, A., Maddison, C.J., Guez, A., Sifre, L., Van Den Driessche, G., Schrittwieser, J., Antonoglou, I., Panneershelvam, V., Lanctot, M. and Dieleman, S., *Mastering the game of Go with deep neural networks and tree search*. Nature, 2016. 529: p. 484.
4. Hinton, G.E., Osindero, S. and Teh, Y.W., *A fast learning algorithm for deep belief nets*. Neural Comput., 2006. 18(7): p. 1527-1554.
5. Zhang, Y., Er, M.J., Venkatesan, R., Wang, N. and Pratama, M., *Sentiment classification using Comprehensive Attention Recurrent models*. in *2016 International Joint Conference on Neural Networks (IJCNN)*. 2016.
6. Yang, L., Zamani H., Zhang Y., Guo J., and Croft W.B., *Neural Matching Models for Question Retrieval and Next Question Prediction in Conversation*. arXiv preprint arXiv:1707.05409, 2017.
7. Zhou, G. and J. Xiangji Huang, *Modeling and Learning Continuous Word Embedding with Metadata for Question Retrieval*. Vol. PP. 2017. 1-1.
8. Severyn, A. and A. Moschitti, *Learning to Rank Short Text Pairs with Convolutional Deep Neural Networks*, in *Proceedings of the 38th International ACM SIGIR Conference on Research and Development in Information Retrieval*. 2015, ACM: Santiago, Chile. p. 373-382.
9. Srba, I. and M. Bielikova, *A comprehensive survey and classification of approaches for community question answering*. ACM Transactions on the Web (TWEB), 2016. 10(3): p. 18.
10. Joty, S., Nakov, P., Màrquez, L. and Jaradat, I., *Cross-language learning with adversarial neural networks*. in *Proceedings of the 21st Conference on Computational Natural Language Learning (CoNLL 2017)*. 2017.
11. Dai, Z. and J. Callan (2019). Deeper text understanding for IR with contextual neural language modeling. *Proceedings of the 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval*.
12. Devlin, J., M.-W. Chang, K. Lee and K. Toutanova (2019). BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*.
13. Pires, T., E. Schlinger and D. Garrette (2019). How Multilingual is Multilingual BERT? *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*.
14. Radlinski, F. and N. Craswell. *A theoretical framework for conversational search*. in *Proceedings of the 2017 Conference on Conference Human Information Interaction and Retrieval*. 2017. ACM.
15. Kiseleva, J., Williams, K., Jiang, J., Hassan Awadallah, A., Crook, A.C., Zitouni, I. and Anastasakos, T., *Understanding user satisfaction with intelligent assistants*. in *Proceedings of the 2016 ACM on Conference on Human Information Interaction and Retrieval*. 2016. ACM.
16. Lowe, R.T., Pow, N., Serban, I.V., Charlin, L., Liu, C.W. and Pineau, J., *Training end-to-end dialogue systems with the ubuntu dialogue corpus*. Dialogue & Discourse, 2017. 8(1): p. 31-65.
17. Gao, J., M. Galley, and L. Li. *Neural approaches to conversational AI*. in *The 41st International ACM SIGIR Conference on Research & Development in Information Retrieval*. 2018. ACM.
18. Chen, H., Liu, X., Yin, D. and Tang, J., *A survey on dialogue systems: Recent advances and new frontiers*. ACM SIGKDD Explorations Newsletter, 2017. 19(2): p. 25-35.

19. Yan, R., Y. Song, and H. Wu. *Learning to respond with deep neural networks for retrieval-based human-computer conversation system*. in *Proceedings of the 39th International ACM SIGIR conference on Research and Development in Information Retrieval*. 2016. ACM.
20. Lowe, R., Pow, N., Serban, I. and Pineau, J., *The Ubuntu Dialogue Corpus: A Large Dataset for Research in Unstructured Multi-Turn Dialogue Systems*. 2015. Association for Computational Linguistics.
21. Yang, L., Qiu, M., Qu, C., Guo, J., Zhang, Y., Croft, W.B., Huang, J. and Chen, H., *Response Ranking with Deep Matching Networks and External Knowledge in Information-seeking Conversation Systems*, in *The 41st International ACM SIGIR Conference on Research & Development in Information Retrieval*. 2018, ACM: Ann Arbor, MI, USA. p. 245-254.
22. Fortuna, P. and S. Nunes, *A survey on automatic detection of hate speech in text*. ACM Computing Surveys (CSUR), 2018. 51(4): p. 85.
23. Nobata, C., Tetreault, J., Thomas, A., Mehdad, Y. and Chang, Y., *Abusive language detection in online user content*. in *Proceedings of the 25th international conference on world wide web*. 2016. International World Wide Web Conferences Steering Committee.
24. Davidson, T., Warmusley, D., Macy, M. and Weber, I., *Automated Hate Speech Detection and the Problem of Offensive Language*. in *ICWSM*. 2017.
25. N. A. Patel and N. Nanavati, "A State of the Art Review on User Behavioral Issues in Online Social Networks," *Recent Advances in Computer Science and Communications*, vol. 16, no. 2, May 2022, doi: 10.2174/2666255815666220513162448
26. R. M. Safari, A. M. Rahmani, and S. H. Alizadeh, "User behavior mining on social media: a systematic literature review," *Multimed Tools Appl*, vol. 78, no. 23, pp. 33747–33804, Dec. 2019, doi: 10.1007/s11042-019-08046-6.
27. Hu, J., Zhang, X., Yang, Y., Liu, Y., & Chen, X. (2020). New doctors ranking system based on vikor method. *International Transactions in Operational Research*, 27(2), 1236–1261.
28. Alarfaj, F., Kruschwitz, U., Hunter, D., & Fox, C. (2012). Finding the right supervisor: expert-finding in a university domain. In *Proceedings of the NAACL HLT 2012 Student Research Workshop* (pp. 1–6).
29. Zhang, X., Cheng, W., Zong, B., Chen, Y., Xu, J., Li, D., & Chen, H. (2020). Temporal context-aware representation learning for question routing. In *Proceedings of the 13th International Conference on Web Search and Data Mining* (pp. 753–761).
30. Liang, S. (2019). Unsupervised semantic generative adversarial networks for expert retrieval. In *The World Wide Web Conference* (pp. 1039–1050).
31. Montandon, J. E., Silva, L. L., & Valente, M. T. (2019). Identifying experts in software libraries and frameworks among github users. In *2019 IEEE/ACM 16th International Conference on Mining Software Repositories (MSR)* (pp. 276–287). IEEE.
32. Liang, S. (2019). Unsupervised semantic generative adversarial networks for expert retrieval. In *The World Wide Web Conference* (pp. 1039–1050).
33. Ambler, S. (2012). *Agile database techniques: Effective strategies for the agile software developer*. John Wiley & Sons.
34. Gharebagh, S. S., Rostami, P., & Neshati, M. (2018). T-shaped mining: A novel approach to talent finding for agile software teams. In *European conference on information retrieval* (pp. 411–423). Springer.

35. Dehghan, M., Biabani, M., & Abin, A. A. (2019). Temporal expert profiling: With an application to t-shaped expert finding. *Information Processing & Management*, 56(3), 1067–1079.
36. Dehghan, M., Rahmani, H. A., Abin, A. A., & Vu, V.-V. (2020). Mining shape of expertise: A novel approach based on convolutional neural network. *Information Processing & Management*, 57(4), 102239.
37. Rostami, P., & Neshati, M. (2019). T-shaped grouping: Expert finding models to agile software teams retrieval. *Expert Systems with Applications*, 118, 231–245.
38. V.-H. Nguyen, K. Sugiyama, P. Nakov, and M.-Y. Kan, “FANG: Leveraging Social Context for Fake News Detection Using Graph Representation,” *Proc. ACM Int. Conf. Inf. Knowl. Manag. CIKM '20*, 2020.
39. C. Yuan, Q. Ma, W. Zhou, J. Han, and S. Hu, “Early Detection of Fake News by Utilizing the Credibility of News, Publishers, and Users based on Weakly Supervised Learning,” *arXiv Prepr. arXiv2012.04233*, pp. 5444–5454, 2020.
40. S. Chandra, P. Mishra, H. Yannakoudakis, M. Nimishakavi, M. Saeidi, and E. Shutova, “Graph-based Modeling of Online Communities for Fake News Detection,” *arXiv Prepr. arXiv2008.06274*, 2020.
41. M. Osmundsen, A. Bor, P. B. Vahlstrup, A. Bechmann, and M. B. Petersen, “Partisan polarization is the primary psychological motivation behind political fake news sharing on Twitter,” *Am. Polit. Sci. Rev.*, vol. 115, no. 3, pp. 999–1015, 2021.
42. K. Shu, A. Sliva, S. Wang, J. Tang, and H. Liu, “Fake News Detection on Social Media : A Data Mining Perspective,” *ACM SIGKDD Explor. Newsletter*, 2017, vol. 19, no. 1, pp. 22–36, 2017.
43. W. L. Hamilton, R. Ying, and J. Leskovec, “Inductive Representation Learning on Large Graphs,” *Neural Inf. Process. Syst.*, 2017.
44. A. Vaswani *et al.*, “Attention Is All You Need,” in *Advances in neural information processing systems*, 2017, pp. 5998–6008.
45. Agrawal, R., Gollapudi, S., Halverson, A. and Ieong, S., 2009, February. Diversifying search results. In *Proceedings of the second ACM international conference on web search and data mining* (pp. 5-14).
46. Narzisi, G., 2008. Yahoo! Clusty-Adding real-time clustering functionality to the Yahoo! web search engine.
47. Zamani, H., Dumais, S., Craswell, N., Bennett, P. and Lueck, G., 2020, April. Generating clarifying questions for information retrieval. In *Proceedings of the web conference 2020* (pp. 418-428).
48. Sekulić, I., Aliannejadi, M. and Crestani, F., 2021. User engagement prediction for clarification in search. In *Advances in Information Retrieval: 43rd European Conference on IR Research, ECIR 2021, Virtual Event, March 28–April 1, 2021, Proceedings, Part I* 43 (pp. 619-633). Springer International Publishing.
49. Ou, W. and Lin, Y., 2020. A clarifying question selection system from NTES_ALONG in Convai3 challenge. *arXiv preprint arXiv:2010.14202*.
50. Zamani, H., Lueck, G., Chen, E., Quispe, R., Luu, F. and Craswell, N., 2020, October. Mimics: A large-scale data collection for search clarification. In *Proceedings of the 29th ACM international conference on information & knowledge management* (pp. 3189-3196).

51. Y. Wang *et al.*, “Enhancing Recommender Systems with Large Language Model Reasoning Graphs,” Aug. 2023, [Online]. Available: <http://arxiv.org/abs/2308.10835>
52. L. Wu, Z. Qiu, Z. Zheng, H. Zhu, and E. Chen, “Exploring Large Language Model for Graph Data Understanding in Online Job Recommendations,” Jul. 2023, [Online]. Available: <http://arxiv.org/abs/2307.05722>