



کاوش داده‌های متنی با استفاده از شبکه‌های عصبی

آزاده شاکری

دانشکده مهندسی برق و کامپیوتر - دانشگاه تهران

چکیده

در سال‌های اخیر، استفاده از شبکه‌های عصبی و به خصوص شبکه‌های عصبی عمیق باعث پیشرفت چشم‌گیری در حوزه‌های مختلف از جمله بینایی ماشین، تشخیص گفتار، و کاربردهای مختلف پردازش زبان طبیعی شده است. با توجه به این پیشرفت‌ها، پژوهش‌های مشابهی برای استفاده از این شبکه‌ها در کاربردهای مختلف بازیابی اطلاعات و کاوش داده‌های متنی آغاز شده است. تمرکز ما در این پژوهش بر روی استفاده از شبکه‌های عصبی و به خصوص شبکه‌های عصبی عمیق در کاربردهای مختلف کاوش داده‌های متنی خواهد بود.

از بین انواع مختلف داده‌ها، داده‌های متنی از اهمیت ویژه‌ای برخوردار هستند و روزانه به شکل‌ها و در قالب‌های مختلف ایجاد می‌شوند. 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. بازیابی پرسش در انجمن‌های پرسش و پاسخ، سامانه‌های مکالمه‌ای، سامانه‌های توصیه‌گر، تشخیص محتوای توهین آمیز و تنفرآمیز، تشخیص اخبار جعلی، و خبره‌یابی از جمله کاربردهایی هستند که در این پژوهش مورد بررسی قرار خواهند گرفت.

کلیدواژه‌ها: مدیریت داده‌های متنی، متن کاوی، یادگیری ماشین، بازیابی اطلاعات، یادگیری عمیق

Text Data Mining using Neural Networks

Abstract

In the recent years, there has been significant improvements in various fields, including machine vision, speech recognition, and natural language processing, due to the advances in neural network models, and especially deep neural networks. Due to these achievements, similar research has begun in the information retrieval and text data mining fields. Our focus in this research will be on the use of neural networks and deep neural networks in various applications of text data mining.

Textual data is of particular importance, among different types of data, and is generated daily in various forms and formats in large volumes, due to the technologies that allow people to easily produce textual data. Examples of textual data include news articles, scholarly articles, books, digital libraries, emails, web pages, content created in Q&A forums, and textual content created on social media. In this research, different applications related to textual data are considered. Question retrieval in Q&A forums, conversation systems, recommendation systems, hate speech detection, fake news detection, and expert finding are among the considered applications.

Keywords: Text data management, Text mining, Machine learning, Information retrieval, Deep learning

مقدمه و بیان مسئله

امروزه با افزایش روز افزون حجم داده‌های برخط رو به رو هستیم. به طور خاص در سال ۲۰۲۰ جهان اساساً تغییر کرد و کووید-۱۹ باعث شد همه جنبه‌های زندگی آنلاین شوند و افراد برای کار، معاشرت، آموزش، و سرگرمی به برنامه‌ها و اینترنت وابسته شوند. در این میان دسترسی سریع و صحیح به منابع مهم و مورد علاقه، یکی از دغدغه‌های استفاده از این 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 سطوح‌های 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. LSTM به پیشرفت ترجمه ماشینی و مدل‌سازی زبانی کمک کرده است. تاکنون شبکه‌های عمیق زیادی معرفی شده‌اند که از جمله پرکاربردترین آن‌ها می‌توان به شبکه عصبی پیچشی^۱، بازگشتی^۲ و شبکه‌های باور عمیق^۳ اشاره کرد. این شبکه‌ها در بازنمایی برداری متن، بازنمایی کلمه، رده‌بندی جمله و موارد دیگر کاربرد دارند [۵].

روش‌هایی که از یادگیری عمیق در کاربردهای داده‌های متنی استفاده می‌کنند، بر خلاف روش‌های سنتی معنی کلمات، گرامر، و نحو را نادیده می‌گیرند و سعی می‌کنند درک متن را بدون داشتن دانشی از زبان مقصد انجام دهند. این روش‌ها برای آموزش معمولاً به حجم زیادی از داده‌های آموزشی نیاز دارند و در بسیاری از کاربردها به طور گسترده موجود هستند. 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 تک‌زبان و هم در محیط دوزبان 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. بکارگیری 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 پرس‌وجو-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

¹ Convolutional Neural Network (CNN)

² Recursive Neural Network (RNN)

³ Deep Belief Network (DBN)

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].

◀ 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 واسط‌های 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 سؤال‌ها و جواب‌های قبلی و همچنین منبع دانش مورد نظر به سؤال 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 سؤال 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 وظیفه‌گرا یا همان چت‌بات‌ها، با هدف پشبرد مکالمه با کاربر با ارائه پاسخ‌هایی منطقی، جملات خبری، جملات احساسی، طرح سؤالات و یا درخواست اطلاعات بیشتر از کاربر طراحی می‌شوند. در این سامانه‌ها معمولاً 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. اخیراً 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 بین زبانی با مدل‌سازی گرافی

تولید محتوای توهین‌آمیز^۴ که در سال‌های اخیر در ارتباطات برخط افزایش یافته است، جرم محسوب می‌شود. عوامل مختلفی در گسترش جملات تنفرآمیز^۵ 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.

⁴ Offensive

⁵ Hate speech

از آنجا که این نوع تعاملات متعصبانه در بستر شبکه‌های اجتماعی می‌تواند برای جامعه مضر باشد، سکوه‌های شبکه اجتماعی می‌توانند از ابزارهای تشخیص و پیش‌گیری در این مورد بهره ببرند.

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 طعن 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،^۶ زبان سوء استفاده،^۷ تبعیض،^۸ دشنام،^۹ و افراطی‌گرایی^{۱۰} 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 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 بین لایه‌ای، 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.

⁶ Cyberbullying

⁷ Abusive language

⁸ Discrimination

⁹ Profanity

¹⁰ Extermism

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 کم‌منابع را افزایش داده‌اند.

کاهش سوگیری در تشخیص سخنان توهین‌آمیز و تنفر آمیز فارسی

در این پژوهش، هدف، انتخاب بهترین مدل زبانی فارسی برای رده‌بندی جملات، بخصوص جملات و سخنان عامیانه، جهت تفکیک و برجسب زنی آن‌ها به کلاس‌های توهین‌آمیز، تنفر آمیز، و یا هیچ کدام است.

این پژوهش شامل سه فاز جمع‌آوری دادگان، به دست آوردن مدل زبانی مطلوب، و سنجش و کاهش سوگیری است. در فاز جمع‌آوری دادگان، هدف به دست آوردن یک مجموعه داده برجسب خورده (دادگان طلایی) جهت آموزش و تست مدل‌های زبانی است. در فاز دوم تحقیق هدف مقایسه بین مدل‌های زبانی مختلف، جهت پیاده‌سازی یک مدل تشخیص کلاس‌های مورد نظر این تحقیق است. با توجه به این که در زبان فارسی، منابع محدودی جهت انجام این کار وجود دارد، و همچنین تحقیقات زیادی در این حوزه تا کنون صورت نگرفته است، این فاز دارای چالش‌ها و نکات قابل تامل فراوانی است. فاز سوم، پیاده‌سازی روش‌های تشخیص سوگیری و همچنین کاهش سوگیری است.

یک روش مبتنی بر بهینه‌سازی برای طراحی خودکار معماری شبکه‌های عصبی همراه با

کاربرد آن در مدل‌های زبانی

یکی از مشکلات مدل‌های مبتنی بر یادگیری عمیق طراحی معماری شبکه است. به خاطر آزادی عمل زیاد در طراحی معماری شبکه‌های عصبی، این مساله تبدیل به یکی از مهم‌ترین بخش‌های یادگیری عمیق شده، و سالانه مقدار زیادی مقاله با هدف طراحی معماری مناسب برای حوزه‌های مختلفی که در آن‌ها از یادگیری عمیق بهره‌بردای میشود به چاپ میرسد. در این پایان‌نامه سعی بر انجام خودکار طراحی شبکه با استفاده از روش‌های مبتنی بر بهینه‌سازی داریم [۲۵].

یکی از مشکلات پیش‌رو در طراحی خودکار معماری شبکه‌های عصبی ایجاد شبکه‌های ناصحیح و بعضا با عملکرد غیر قابل قبول توسط این روش‌ها است. به علاوه، در حالی که انتظار می‌رود انجام عملیات جستجو به صورت طولانی مدت سبب یافتن معماری‌های با عملکرد بهتر بشود، مشاهده شده‌است که انجام طولانی مدت عملیات جستجو منجر به ایجاد معماری‌های با عملکرد پایین و ساختار نامناسب می‌شود [۲۵-۲۶]. در این پایان‌نامه، به صورت ریاضی و تجربی، این مشاهده به ساختار چند لایه‌ای معماری شبکه‌های عصبی عمیق نسبت داده می‌شود.

به خاطر ساختار چندلایه‌ای، معماری مناسب برای لایه‌های بالاتر عموماً دارای تعداد زیادی عملگر غیر پارامتری (مانند لایه‌های Pooling) می‌باشند و معماری مناسب برای لایه‌های پایین‌تر عموماً دارای عملگرهای پارامتری (مانند لایه‌های Convolution) است. به خاطر این اختلاف، استفاده از روش‌های بهینه‌سازی مبتنی بر گرادیان سبب ایجاد معماری‌های با عملکرد ضعیف می‌شود. در این پایان‌نامه، با معرفی یک عبارت منظم‌سازی کننده در تابع هدف، از این اتفاق جلوگیری می‌کنیم. در این عبارت، شبکه به نزدیک کردن معماری بهینه برای لایه‌های مختلف شبکه تشویق می‌شود.

◀ خبره‌یابی

مسئله‌ی خبره‌یابی یکی از مسائل شناخته شده در حوزه‌ی بازیابی اطلاعات است که تا به الان در زمینه‌های مختلفی همچون یافتن افراد مناسب برای داوری مقالات [۲۷] پیشنهاد دکتر مناسب به بیماران [۲۸]، یافتن استاد راهنما برای دانشجویان [۲۹] و مسیربازی سوال [۳۰] مورد استفاده قرار گرفته است. اما یکی از مهم‌ترین کاربردهای آن انتخاب افراد مناسب از بین مجموعه‌ای از داوطلبان برای یک موقعیت شغلی [۳۱] داده شده است. در اینجا به‌طور کلی هدف این است که با در اختیار داشتن لیستی از حوزه‌های مهارتی مورد نیاز یک موقعیت شغلی، یک رتبه‌بندی از داوطلبان بر اساس میزان دانششان در حوزه‌های مهارتی ذکر شده تولید شود. تا به حال تحقیقات زیادی در این زمینه انجام شده است [۳۲-۳۳]، اما بسیاری از آنها قادر به برآورده کردن نیازهای اصلی شرکت‌ها نیستند و بدون توجه به این نیازها 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، به‌طوریکه هر یک در یک حوزه‌ی مورد نیاز تیم خبره و در سایر حوزه‌های مورد نیاز دانش عمومی دارد [۳۴]. 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 [۳۵-۳۷]. 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 دو مدل احتمالاتی مولد که بر اساس تعداد اسناد داوطلبان، دانش موردنیاز آنها را تخمین می‌زنند اقدام به بازیابی تیم‌های چابک نمودند. با این حال، آنها هیچ توجه‌ای به محتوی اسناد داوطلبان نداشتند و دانش داوطلبان را بدون توجه به محتوی اسنادشان تخمین می‌زدند.

ما در اینجا قصد داریم که مدل‌های یادگیری عمیقی پیشنهاد دهیم که قادر باشند از محتوی اسناد داوطلبان استفاده کنند تا مناسب‌ترین گروه از داوطلبان را برای تشکیل یک تیم چابک پیشنهاد دهند. به‌طور دقیق‌تر ایده‌ی ما به این صورت است که ابتدا می‌خواهیم با استفاده از یک رویکرد مبتنی بر رتبه‌بندی مجدد که از مدل یادگیری عمیق پیشنهادی ما کمک می‌گیرد داوطلبان منتخب را بازیابی کنیم. سپس با استفاده از یک مدل برنامه‌ریزی خطی عدد صحیح، اعضای مناسب تیم چابک را از بین داوطلبان منتخب انتخاب نماییم.

برای ایجاد مدل یادگیری عمیق ذکر شده، ما در قدم اول قصد داریم که از فرآیند تغییر توزیع تاپیک‌های اسناد داوطلبان در گذر زمان استفاده کنیم تا بتوانیم پروفایلی برایشان ایجاد کنیم، سپس با مقایسه‌ی پروفایل ایجاد شده و بازنمایی‌های حوزه‌های مهارتی مورد نیاز تیم، دانش داوطلبان برای امکان حضور در یک تیم چابک را تخمین بزنیم. در قدم دوم ما قصد داریم که از مکانیزم مبتنی بر توجه استفاده کنیم تا با درک دقیق موضوع اسناد داوطلبان، مناسب‌ترین پروفایل را برایشان ایجاد کنیم، سپس از مقایسه‌ی پروفایل آنها با بازنمایی‌های حوزه‌های مهارتی مورد نیاز تیم، دانش آنها را همچون قدم اول تخمین بزنیم.

◀ تشخیص اخبار جعلی در شبکه‌های اجتماعی مبتنی بر سوگیری‌های شناختی

در سال‌های اخیر شبکه‌های اجتماعی به محبوبیت زیادی میان عموم جامعه دست یافته‌اند. دسترسی آسان و انتشار سریع سبب شده است که این شبکه‌ها به عنوان منبع دریافت انواع اخبار برای کاربران قرار گیرند. از طرفی با توجه به اینکه در شبکه اجتماعی هر کاربر می‌تواند بدون محدودیت اطلاعاتی را منتشر کند انتشار خبرهای جعلی در این شبکه‌ها در حال افزایش است. بنابراین تشخیص خودکار اخبار جعلی در شبکه‌های اجتماعی بسیار مورد توجه قرار گرفته است. با توجه به شباهت بسیار میان متن اخبار جعلی و حقیقی، روش‌های موجود علاوه بر متن خبر، از بافتار شبکه‌های اجتماعی نظیر کاربران و نحوه انتشار، جهت صحت‌سنجی اخبار استفاده کرده‌اند.

اعتبار خبر منتشر شده در شبکه اجتماعی وابسته به اعتبار کاربر منتشرکننده‌ی خبر و همچنین مورد اعتماد بودن منبع خبر می‌باشد. هر چند اعتبار منبع خبر و همچنین کاربران منتشر کننده خبر جهت صحت‌سنجی خبر در پژوهش‌های پیشین [۳۸-۴۰] بسیار مورد توجه قرار گرفته است ارتباط میان موضوع و سوگیری سیاسی خبر، سوگیری اجتماعی-شناختی کاربران منتشر کننده و سوگیری سیاسی منبع خبر در صحت‌سنجی اخبار منتشر شده نادیده گرفته شده است. در این پژوهش‌ها با توجه به سابقه کاربر یا منبع خبر در انتشار اخبار جعلی، اعتبار آن‌ها مشخص می‌شود، هر چند اعتبار کاربران و منابع خبری با توجه به سوگیری‌های شناختی آن‌ها در موضوعات مختلف متفاوت می‌باشد. طبق نظریه سوگیری تاییدی^۱ که در حوزه روانشناسی مطرح می‌شود، افراد مفاهیمی را که مطابق با اعتقادات و باورهایشان است باور می‌کنند. طبق این نظریه، افراد تمایل دارند اطلاعات نزدیک به دیدگاهشان را دریافت و مصرف کنند. این دیدگاه در تصمیم‌گیری آن‌ها برای انتشار یک مطلب در شبکه‌های اجتماعی تاثیر می‌گذارد [۴۱]. بنابراین اعتبار اخبار باز نشر شده توسط کاربران در شبکه‌های اجتماعی و همچنین اخبار منتشر شده توسط منابع خبری، با توجه به دیدگاه و همچنین میزان سوگیری‌های آن‌ها در موضوع خبر متفاوت است. به عنوان مثال کاربری که مخالف یک جناح سیاسی است، اخبار ضد این جناح را با دقت کمتری نسبت به صحت خبر و صرف اینکه خبر مطابق با سوگیری وی می‌باشد باز نشر می‌کند.

علاوه بر سوگیری‌های شناختی کاربر، ساختار شبکه و اثر نفوذ سایر کاربران و انجمن‌هایی که کاربر در شبکه اجتماعی به آن تعلق دارد در دیدگاه کاربر و در نتیجه رفتار انتشاری کاربران در شبکه‌های اجتماعی تاثیر می‌گذارد. سوگیری تاییدی علاوه بر انتشار اخبار توسط کاربران، روی صفحات و افرادی که یک کاربر در شبکه دنبال می‌کند نیز تاثیر می‌گذارد و باعث ایجاد انجمن‌هایی از افراد همگون در شبکه‌های اجتماعی می‌شود. در این انجمن‌ها هر فرد اخبار و موضوعاتی را دریافت و منتشر می‌کند که مطابق با اعتقادات و باورهای فرد باشد که این اثر اتاق پژواک^۲ نام دارد. اثر اتاق پژواک باعث ایجاد انجمن‌های متعصبانه در شبکه می‌شود که نسبت به موضوعات مختلف به صورت عمد یا غیر عمد حالت خصمانه و متعصبانه پیدا می‌کنند و سوگیری اجتماعی-شناختی کاربران را تقویت می‌کنند که سبب افزایش انتشار اخبار جعلی در

¹ Confirmation Bias

1

¹ Echo –Chamber Effect

2

این انجمن‌ها می‌شود [۴۲]. سوگیری اجتماعی-شناختی شامل دیدگاه‌هایی است که بین اعضای این انجمن‌ها مشترک است.

در این پژوهش هدف ما ایجاد مدل شبکه عمیقی است که با توجه به ارتباط بین موضوع و موجودیت‌های خبرهای منتشر شده، سوگیری اجتماعی-شناختی کاربران انتشار دهنده‌ی این اخبار و همچنین سوگیری سیاسی منابع خبرها، اعتبار کاربران و منبع خبر را نسبت به موضوع خبر مدل کرده و جهت صحت‌سنجی خبر منتشر شده به کار گیریم. بدین منظور ابتدا خبر، کاربران و منبع خبر را با استفاده از شبکه‌های عصبی گرافی [۴۲] به گونه‌ای بازنمایی می‌کنیم که بردارهای بازنمایی حاصل به ترتیب بیانگر موضوع و سوگیری سیاسی خبر (با استفاده از موجودیت‌های خبر)، سوگیری اجتماعی-شناختی کاربر (با استفاده از ارتباطات اجتماعی قوی اطراف کاربر و متن اخبار منتشر شده توسط آن‌ها) و دیدگاه و سوگیری سیاسی منبع خبر (با بهره‌گیری از مخاطبان مشترک بین منابع خبر) باشد. سپس با الهام از مکانیزم توجه در شبکه‌های عمیق [۴۳] ماژولی طراحی کرده که تعامل میان سه بردار خبر، کاربر و منبع خبر را به گونه‌ای مدل می‌کند که اعتبار کاربر و منبع خبر نسبت به موضوع خبر به صورت ضمنی مدل شده و جهت صحت‌سنجی اخبار به کار گرفته می‌شود.

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

بهبود در سیستم‌های مکالمه

امروزه سیستم‌های مکالمه کاربرد زیادی پیدا کرده اند. یکی از چالش‌هایی که در این سیستم‌ها وجود دارد، عدم فهم صحیح هدف کاربر از ورودی اعلام شده است. این مساله باعث می‌شود که سیستم قابلیت پاسخگویی به کاربر را نداشته باشد. در این راستا لازم است که با مطرح نمودن سوال شفاف‌سازی مناسب، تلاش شود تا ابهام موجود برطرف گردد [۴۵-۴۷].

در این راستا نیاز است تا دو گام اصلی انجام شود:

- شناسایی سطح ابهام در ورودی کاربر
- پرسش سوال شفاف‌سازی مناسب براساس ورودی کاربر و سابقه مکالمه

در راستای گام اول، باید ماژولی طراحی و پیاده‌سازی شود که پس از دریافت آخرین اظهار کاربر به همراه سابقه مکالمه، سطح ابهام را در قالب عددی صحیح در بازه [۴۵]. به عنوان خروجی اعلام کند. برای این دستیابی به این هدف، لازم است تا نمایش معنایی مناسبی از اطلاعات ورودی ایجاد و کلاسه‌بند مناسب پیاده‌سازی شود.

در راستای گام دوم، نیاز داریم تا با در نظر گرفتن بانک سوالات شفاف‌سازی، پس از ایجاد نمایش مناسب از اطلاعات ورودی و سوالات موجود، بهترین پرسش شناسایی و از کاربر پرسیده شود. در ادامه، از آنجایی که امکان دارد رفع ابهام در بیش از یک اظهار حاصل شود، لازم است تا مجدداً پس از اطمینان از وجود ابهام، سوال جدید پرسیده شود. برای مطرح

نمودن سوال جدید، علاوه بر سابقه مکالمه، لازم است به صورت خاص از اطلاعات موجود در پاسخ کاربر بهره مند شویم. همچنین لازم است وزن اهمیت پاسخ کاربر را نسبت به اطلاعات موجود در سابقه مکالمه بیشتر کنیم.

➤ تشخیص اخبار جعلی بین زبانی بر اساس محتوا در زبان های کم منبع

در دهه‌ی اخیر با توسعه‌ی سریع فناوری و ظهور اینترنت، رسانه‌های اجتماعی آنلاین روند انتشار اخبار را از نظر سرعت انتشار و موانع دسترسی مخاطبان به اطلاعات، به طور قابل توجهی تغییر داده‌اند. بررسی‌ها نشان داده‌است که ۹۳,۳۳٪ از ۴,۸ میلیارد کاربر جهانی اینترنت در شبکه‌های اجتماعی آنلاین هستند [۴۸].

در گذشته اخبار تنها از طریق رسانه‌های خبری معمول مانند رسانه‌های تلویزیونی و روزنامه‌ها به اشتراک گذاشته می‌شد، اما امروزه شبکه‌های اجتماعی آنلاین به عنوان بستر اصلی اشتراک‌گذاری اخبار و نظرات کاربران مورد استفاده قرار می‌گیرد. علت این امر این است که توزیع محتوا در این رسانه‌ها در مقایسه با رسانه‌های سنتی، اغلب سریعتر و با هزینه‌ی کمتری است [۴۹]. در دسترس بودن و ناشناس بودن شبکه‌های اجتماعی آنلاین، اشتراک‌گذاری و تبادل اطلاعات را برای کاربران راحت‌تر کرده‌است، اما از طرفی سبب فراهم شدن بستری برای انتشار اخبار جعلی نیز شده‌است. به عنوان مثال در انتخابات ۲۰۱۶ ریاست جمهوری آمریکا، انتشار اطلاعات اشتباه از طریق رسانه‌های اجتماعی آنلاین منجر به تغییر رای جامعه و تغییر نتیجه‌ی انتخابات شد [۵۰-۵۱]. این اتفاقات و آسیب‌های ناشی از آن‌ها، انگیزه‌ای برای تحقیقات آینده بر روی تشخیص اخبار جعلی ایجاد کرد.

در طول زمان، بارها تاثیر انتشار اخبار جعلی بر پدیده‌های مختلف مانند پدیده‌های سیاسی، اجتماعی و اقتصادی را مشاهده کرده‌ایم. از این رو بسیاری از نهادهای دولتی، سازمان‌ها و شرکت‌ها به دنبال راه‌حلی برای حل این مشکل هستند. گسترش اخبار جعلی نه تنها می‌تواند باعث گمراهی افراد و اقدامات گمراه کننده شود؛ بلکه میزان اعتماد کاربران به اخبار را به طور قابل توجهی کاهش می‌دهد. امروزه اکثر تحقیقات در زمینه‌ی پردازش زبان‌های طبیعی که تشخیص اخبار جعلی نیز جزو آن‌ها محسوب می‌شود بر روی زبان‌هایی با منابع زیاد و برچسب خورده مانند زبان انگلیسی انجام می‌شود، در نتیجه تعداد زیادی از زبان‌های دیگر در سراسر دنیا نادیده گرفته می‌شوند. استفاده از سیستم‌هایی بر پایه‌ی یادگیری ماشین و یا شبکه‌های عصبی نتایج تحقیقات در زمینه‌ی پردازش زبان را به صورت چشمگیری بهبود داده است، اما این سیستم‌ها تنها بر روی زبان‌هایی با داده‌های زیاد خوب عمل می‌کنند. در نتیجه پردازش زبان طبیعی بر روی زبان‌هایی با منابع کم به عنوان یک چالش مهم مطرح شده است [۵۲]. از آنجایی که ساخت دستی مجموعه داده‌های برچسب خورده برای تمام زبان‌های کم‌منبع دنیا ناکارآمد است، می‌توان از روش‌های یادگیری انتقال بین زبانی برای رسیدن به این هدف استفاده کرد [۵۳].

کارهای محدودی در زمینه‌ی انتقال دانش بین‌زبانی در تشخیص اخبار جعلی بر روی زبان‌های کم‌منبع انجام شده است. مطالعه‌ی [۵۴] روشی برای تشخیص اخبار جعلی با استفاده از ویژگی‌های دست‌ساز مستقل از متن در سه زبان انگلیسی اسپانیایی و پرتغالی ارائه می‌دهد. مقاله‌ی [۵۵] به تشخیص شایعه در زبان‌های انگلیسی هندی و بنگالی با استفاده از ویژگی‌های شبکه و کاربر و ویژگی‌های متنی با استفاده از تنظیم دقیق Mbert پرداخته است. یکی دیگر از روش‌های

مورد بحث در انتقال دانش بین‌زبانی استفاده از شبکه‌های متخاصم است. پژوهش [۵۶] به استفاده از شبکه‌های متخاصم زبانی در جهت رده‌بندی احساسات بین‌زبانی با استفاده از تولید ویژگی‌های مستقل از زبان پرداخته است. همانطور که می‌بینیم بیشتر این مقالات یا تنها از ویژگی‌های دست‌ساز استفاده کرده‌اند که نمی‌تواند ویژگی‌های پیچیده‌تر در تشخیص اخبار جعلی را فراهم کند یا به استفاده از بردارهای از پیش‌آموزش دیده‌ی چندزبانه اکتفا کرده‌اند. یکی از چالش‌هایی که در مورد بردارهای جاسازی از پیش‌آموزش داده شده‌ی چندزبانه وجود دارد این است که برخی از این بردارها مانند MBert [۵۶] حتی بعد از تنظیم دقیق هم خیلی از نظر زبانی بی‌طرف نیستند و پدیده‌های مشابه به طور یکسان در زبان‌ها نمایش داده نمی‌شوند، در حقیقت این بردارها برای زبان‌هایی با منابع کم به اندازه‌ی زبان‌های پر استفاده خوب عمل نمی‌کنند [۵۷]. از طرفی یکی از مباحثی که در این پژوهش‌ها در نظر گرفته نشده است وجود دامنه‌های مختلف اخبار جعلی در زبان‌های مختلف است. برای مثال پژوهش [۵۸]. به تولید بردارهای جاسازی مستقل از دامنه برای تشخیص اخبار جعلی در یک زبان پرداخته است که به نتایج خوبی نیز دست پیدا کرده است. هدف در این پژوهش بهبود دقت مدل‌های تشخیص اخبار جعلی در زبان‌های کم‌منبع با استفاده از بهبود بردارهای جاسازی با وارد کردن یکسری از ویژگی‌های مرتبط با اخبار جعلی (برای مثال دامنه) یا بهبود در جهت مستقل از زبان شدن این بردارها و ترکیب آن با ویژگی‌های دست‌ساز متنی است.

◀ بازیابی پرسش شفاف‌ساز در سیستم‌های جستجوی مکالمه‌ای

موتورهای جستجو اهمیتی 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 [۵۹] یا دسته‌بندی نتایج [۶۰] -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 [۶۱] است. در این روش موتور جستجو با ارائه پرسشی در پاسخ به پرس‌وجوی مبهم و مشخص کردن تعدادی پاسخ کاندیدا 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 [۶۲] 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 [۶۳] برای پرس و جوی 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 پرس و جوی مطرح شده، و استفاده از آن در کنار پرس و جو، برای تشخیص ابهام و همچنین به منظور بازیابی پرسش شفاف‌ساز مناسب، به نتایجی بهتر از آنچه در مطالعات قبلی [۶۴] به آن دست یافته‌اند برسیم. در این مسیر از مجموعه داده‌ای با نام MIMICS که از داده‌های پرسش شفاف‌ساز موجود در موتور جستجو بینگ جمع‌آوری شده است، استفاده کنیم.

سیستم‌های توصیه‌گر مکالمه‌ای

با گسترش وبسایت‌های تجارت الکترونیک، محصولات و خدمات ارائه شده در فضای وب به سرعت در حال افزایش است. تنوع بسیار بالای محصولات موجود در این فضا سبب تحت تاثیر قرار دادن کاربران می‌گردد و گاهی به جای سودمندی برای آنان، ممکن است تصمیم‌گیری را برای کاربران سخت نماید، در نتیجه سبب انتخاب‌های نامناسب گردد. با توجه به اینکه سیستم‌های توصیه‌گر یک ابزار حیاتی در جهت بهبود تجارب کاربران و همچنین ارتقای نرخ فروش و ارائه خدمات در بسیاری از وبسایت‌های آنلاین محسوب می‌گردند [۶۵]، در سال‌های اخیر علاقه به این سیستم‌ها به شکل چشم‌گیری افزایش یافته است. سیستم‌های توصیه‌گر براساس نوع داده‌های ورودی به سه دسته‌ی اصلی: فیلتر مشارکتی^۱، مبتنی بر محتوا^۲ و ترکیبی^۳ تقسیم می‌شوند [۶۶].

سیستم‌های توصیه‌گر بیان شده، سیستم‌های توصیه‌گر ایستا نامیده می‌شوند [۶۶] و در سال‌های اخیر پژوهش‌های بسیاری در این حوزه صورت پذیرفته است. علی‌رغم موفقیت‌های پژوهش‌های صورت گرفته در این حوزه اما این سیستم‌ها دارای محدودیت‌هایی مختلفی هستند، که می‌توان به موارد زیر اشاره کرد:

- سیستم‌های توصیه‌گر ایستا براساس ترجیحات پیشین کاربران به آن‌ها توصیه‌هایی ارائه می‌دهند و نمی‌توانند براساس علایق فعلی کاربران، پیشنهادهای ارائه نمایند.
- در سیستم‌های توصیه‌گر ایستا دلیل توصیه‌ی ارائه شده به هر کاربر بیان نمی‌گردد [۶۶]، در واقع دارای محدودیت عدم تشریح‌پذیری^۴ هستند.
- توصیه‌های ارائه شده توسط سیستم‌های ایستا معمولاً براساس خواسته‌های ضمنی کاربران بیان شده است و این سیستم‌ها به منظور استخراج خواسته‌های صریح کاربران با کاربران به شکل مستقیم تعامل ندارند [۶۷].

1 Collaborative filtering	3
1 Content based	4
1 Hybrid	5
1 Explainability	6

- از دیگر محدودیت‌های این سیستم‌ها، مشکل شروع سرد است، در این حالت چون از کاربران جدید اطلاعاتی وجود ندارد و با کاربر تعاملی شکل نگرفته است، بنابراین اقلام توصیه شده از دقت کافی برخوردار نیستند.
 - اگر چه سیستم‌های توصیه‌گر ایستا تلاش‌های بسیاری در جهت حل مشکل سربار^۷ اطلاعات برای کاربران در وبسایت‌های تجارت الکترونیک انجام داده است، اما هنوز هم تعداد زیادی انتخاب را برای کاربران باقی می‌گذارد.
 - در سیستم‌های توصیه‌گر ایستا توصیه‌های انجام شده معمولاً فاقد تنوع و تازگی هستند و اقلام پرتعدادتر همیشه نسبت به اقلام کم‌تعدادتر به کاربران توصیه می‌شوند که این موضوع بایاس محبوبیت^۸ نامیده می‌شود.
 - بسیاری از سیستم‌های توصیه‌گر ایستا با این فرض ایجاد شده است که کاربر از قبل ترجیحات خود را می‌داند، اما این موضوع لزوماً درست نیست. کاربران ممکن است از ترجیحات خود حین فرایند تصمیم‌گیری و زمانی که از گزینه‌های موجود آگاه می‌شوند، مطلع گردند.
 - بعضی از اقلام دارای ویژگی‌های فراوانی هستند، که تمامی این ویژگی‌ها باید با توجه به خواسته‌های هر کاربر شخصی‌سازی گردد. مانند تلفن همراه که دارای ویژگی‌های فراوانی مانند اندازه صفحه نمایش، کیفیت دوربین، سیستم‌عامل، نوع شبکه و موارد دیگر هستند. سیستم‌های توصیه‌گر ایستا در توصیه این اقلام مشکل دارند.
- به لطف پیشرفت‌های صورت گرفته در حوزه پردازش زبان طبیعی، تاثیر مثبت چت‌بات‌های تجاری در افزایش نرخ خرید و فروش کالا و با توجه به محدودیت‌های بیان شده برای سیستم‌های توصیه‌گر ایستا، سیستم‌های توصیه‌گر مکالمه‌ای^۹ در سال‌های اخیر مورد توجه قرار گرفته‌اند. یک سیستم توصیه‌گر مکالمه‌ای نرم‌افزاری است که می‌تواند به زبان طبیعی با کاربران تعامل داشته باشد و بازخورد صریح کاربران را به منظور درک علایق پویای آن‌ها دریافت کند و در مقابل توصیه‌هایی به هر کاربر ارائه بدهد.
- در سال‌های اخیر پژوهش‌های مختلفی به منظور ارائه روش‌هایی برای سیستم‌های توصیه‌گر مکالمه‌ای ارائه شده است. در این روش‌ها بیان شده است که ارائه توصیه صرفاً براساس چندین مکالمه‌ی کوتاه بین سیستم و کاربر، پیچیده و دشوار است، به همین منظور در مقالاتی انواع مختلف دانش‌های خارجی ساختاریافته و غیرساختاریافته مانند گراف دانش در سطح اقلام، گراف دانش در سطح کلمات، نقدهای بیان شده برای اقلام و حتی داده‌های تعاملات گذشته‌ی کاربران در یک سیستم به داده‌های مکالمه‌ای اضافه گردید و سبب افزایش کیفیت توصیه‌ها و گفتگوهای تولید شده توسط سیستم شد [۶۸].
- پس از بررسی کلی روش‌های مطرح شده در این حوزه اکنون به بررسی مسائل باز این حوزه و همچنین ایده‌هایی که در ادامه‌ی کار می‌توان به آن‌ها پرداخت اشاره می‌شود.

1 Overload 7
 1 Popularity bias 8
 1 Conversational Recommender System (CRS)

- استفاده از دانش خارجی:

همانطور که بیان گردید استفاده از دانش خارجی در بهبود کیفیت سیستم‌ها اثر گذار است، بر همین اساس در ادامه‌ی کار می‌توان از دانش‌های خارجی دیگر مانند اطلاعات دموگرافیک کاربران استفاده کرد.

- مسئله‌ی بایاس محبوبیت:

از دیگر مسائلی که می‌توان به آن پرداخت، مشکل بایاس محبوبیت است، گر چه سیستم‌های توصیه‌گر مکالمه‌ای به نسبت سیستم‌های توصیه‌گر ایستا موارد تکراری و پرتعداد را کمتر توصیه می‌کنند، اما هنوز هم این مشکل در بسیاری از سیستم‌های توصیه‌گر مکالمه‌ای وجود دارد و می‌توان راهکارهایی را در جهت کاهش این مشکل ارائه کرد.

- بهبود روش‌های شبیه سازی کاربران:

همانطور که بیان گردید، داده‌های شبیه سازی کاربران به دلیل مشکل کمبود داده واقعی برای اهداف متفاوتی مانند تنظیم کردن پارامترهای اولیه در مدل‌ها استفاده می‌شوند. اما این داده‌های شبیه سازی شده دارای مشکلات مختلفی هستند. بنابراین به عنوان یکی از مسائل باز این حوزه می‌توان به بهبود نحوه‌ی شبیه‌سازی کاربران اشاره کرد.

◀ تشخیص سخنان تنفرآمیز فارسی با در نظر گرفتن ساختار اجتماعی کاربران

افزایش محتوای تولید شده توسط کاربران در اینترنت، به ویژه در شبکه‌های اجتماعی، میزان سخنان تنفرآمیز نیز به طور پیوسته افزایش یافته است. انتشار این قبیل محتوا در شبکه‌های اجتماعی آثار مخرب فرهنگی و اجتماعی دارد که در صورت عدم رسیدگی به موقع به این محتوا می‌تواند منجر به وقوع اتفاقات جبران ناپذیری شود [۶۹-۷۰]. حتی برخی از مقالات از این موارد به عنوان محرک‌های اقدام به خودکشی، افکار خودکشی و مشکلات سلامت روان مورد بررسی اشاره شده‌اند. بر همین اساس در ایالات متحده آمریکا، سخنان تنفرآمیز تحت پوشش اصلاحیه آزادی بیان قانون اساسی است [۷۱] و همچنین قانون رفتار اتحادیه اروپا رسانه‌های اجتماعی را ملزم به شناسایی و حذف سخنان تنفرآمیز در کمتر از ۲۴ ساعت از انتشار آن‌ها می‌کند [۷۲].

در پژوهش‌هایی که انجام شده است، تعاریف مختلفی برای سخنان تنفرآمیز ذکر شده است. یکی از تعاریف مرسوم آن بیان می‌کند که سخنان تنفرآمیز به هر گونه ارتباطی گفته می‌شود که فرد یا گروهی را بر اساس برخی ویژگی‌ها مانند نژاد، رنگ، قومیت، جنسیت، گرایش جنسی، ملیت، مذهب یا سایر ویژگی‌های دیگر تحقیر شوند. به طور کلی تشخیص سخنان تنفرآمیز عبارت است از تشخیص اینکه آیا ارتباطاتی حاوی موارد تنفرآمیز هستند و یا اینکه باعث تشویق خشونت نسبت به یک فرد یا گروهی از افراد می‌شوند یا خیر [۷۳-۷۴]. تا به حال، تحقیق‌ها و مطالعه‌های بسیاری درباره انتشار سخنان تنفرآمیز صورت گرفته است، اما در بیشتر رویکردها جنبه تعامل کاربران را در نظر نمی‌گیرند. تصمیم‌گیری در مورد اینکه آیا برخی از گفته‌ها بیانگر سخنان تنفرآمیز هستند یا خیر، مسئله‌ای دشوار است. تشخیص سخنان تنفرآمیز، تنها با نگاه کردن به کلمات کلیدی قابل حل نیست، این که آیا یک پیام تنفرآمیز است یا خیر، به اطلاعات زمینه‌ای که داده می‌شود بستگی دارد و از آنجا که طبق پژوهش‌ها معمولاً افرادی که باهم در ارتباط هستند، رفتار مشابهی را انجام از خود نشان

می‌دهند و در مطالعات این حوزه به پدیده، اثرات اتاق صدای تکراری اشاره می‌کنند که بر طبق آن افراد از اطلاعات و نظراتی که با باورهای و تعصبات موجودشان سازگاری دارد، استفاده می‌کنند لذا استفاده از ساختارها اجتماعی در این حوزه می‌تواند بسیار موثر باشد [۷۵-۷۹].

تشخیص سخنان تنفرآمیز یک وظیفه پرچالش است. طبق پژوهش‌های انجام شده، از مهم‌ترین چالش‌ها در این حوزه ظرافت‌های زبانی و بیان غیرمستقیم این عبارات و محدودیت‌های در دسترس بودن داده‌ها برای آموزش و آزمایش مدل‌های خودکار برای تشخیص سخنان تنفرآمیز است [۷۳].

رسانه‌های اجتماعی، قانونی برای سخنان تنفرآمیز هستند. از طرفی با توجه به اینکه سیاست‌های استفاده و توزیع داده بسیار سخت‌گیرانه‌ای دارند، برخی از مجموعه داده‌های تنظیم‌شده به صورت عمومی منتشر نشده است. از سوی دیگر، با توجه به اینکه کمتر پژوهشی در زبان فارسی انجام شده است و مجموعه داده‌ای که به زبان فارسی باشد و شامل ساختار اجتماعی کاربران و متن ارتباط بین کاربران شود، تولید نشده است، ایجاد مجموعه داده با در نظر گرفتن ساختار اجتماعی بین کاربران و ایجاد مدلی مبتنی بر آن از اهمیت زیادی برخوردار است.

ما در این پژوهش قصد داریم در گام نخست از داده‌های شبکه اجتماعی پاساژ استفاده کنیم و مجموعه داده‌ای با در نظر گرفتن ساختار اجتماعی بین کاربران ایجاد نماییم، با توجه به اینکه بیان کردیم استفاده از ساختارهای اجتماعی بین کاربران می‌تواند به بهبود عملکرد این وظیفه کمک کند، براساس مجموعه داده‌ای که استخراج کردیم از ساختارهای اجتماعی استفاده خواهیم کرد تا کاربران و نحوه تعامل آن‌ها را به صورت مناسب مدل کنیم و با استفاده از روش‌های مناسب و استفاده از تحلیل شبکه‌های اجتماعی، همسایگی، زمینه و الگوهای تعاملی یک کاربر را برای آن تعریف خواهیم کرد تا در طبقه‌بندی کاربرانی که سخنان تنفرآمیز تولید می‌کنند استفاده کنیم تا در نهایت مدلی به منظور تشخیص خودکار سخنان تنفرآمیز ایجاد کنیم.

جمع بندی

در این پژوهش، تمرکز بر روی استفاده از شبکه‌های عصبی در چند کاربرد مختلف کاوش داده‌های متنی و بازیابی اطلاعات است. به طور خاص در این پژوهش، سامانه‌های مکالمه‌ای جویای اطلاعات، سامانه‌های توصیه‌گر، تشخیص محتوای توهین آمیز و تنفرآمیز، تشخیص اخبار جعلی، و خبره‌یابی خواهد بود.

² Homophily	0
² Echo Chamber Effect	1

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. Thomas Elsken, Jan Hendrik Metzen, and Frank Hutter. 2019. Neural Architecture Search: A Survey. *J. Mach. Learn. Res.* 20, (2019), 55:1–55:21.
26. Xiangxiang Chu, Tianbao Zhou, Bo Zhang, and Jixiang Li. Fair DARTS: eliminating unfair advantages in differentiable architecture search. In Andrea Vedaldi, Horst Bischof, Thomas Brox, and Jan-Michael Frahm, editors, *Computer Vision - ECCV 2020 - 16th European*
27. Mimno, D., & McCallum, A. (2007). Expertise modeling for matching papers with reviewers. In *Proceedings of the 13th ACM SIGKDD international conference on knowledge discovery and data mining* (pp. 500–509).
28. 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.
29. 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).
30. 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).
31. Liang, S. (2019). Unsupervised semantic generative adversarial networks for expert retrieval. In *The World Wide Web Conference* (pp. 1039–1050).
32. 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.

33. Liang, S. (2019). Unsupervised semantic generative adversarial networks for expert retrieval. In *The World Wide Web Conference* (pp. 1039–1050).
34. Ambler, S. (2012). *Agile database techniques: Effective strategies for the agile software developer*. John Wiley & Sons.
35. 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.
36. 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.
37. 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.
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. Aliannejadi, M., Zamani, H., Crestani, F., & Croft, W. B. (2019, July). Asking clarifying questions in open-domain information-seeking conversations. In *Proceedings of the 42nd international acm sigir conference on research and development in information retrieval* (pp. 475-484).
46. Bi, K., Ai, Q., & Croft, W. B. (2021, July). Asking Clarifying Questions Based on Negative Feedback in Conversational Search. In *Proceedings of the 2021 ACM SIGIR International Conference on Theory of Information Retrieval* (pp. 157-166).
47. Rao, S., & Daumé III, H. (2018). Learning to ask good questions: Ranking clarification questions using neural expected value of perfect information. *arXiv preprint arXiv:1805.04655*.
48. B. Dean, “How Many People Use Social Media in 2021?,” *Backlinko*, 2021. <https://backlinko.com/social-media-users>.

49. X. Zhou and R. Zafarani, "Fake News: Survey of Research, Detection Methods, and Opportunities," *CoRR*, vol. abs/1812.0, 2018, [Online]. Available: <http://arxiv.org/abs/1812.00315>.
50. H. Allcott and M. Gentzkow, "Social media and fake news in the 2016 election," *J. Econ. Perspect.*, vol. 31, no. 2, pp. 211–236, 2017, doi: 10.1257/jep.31.2.211.
51. K. Rapoza, "Can 'Fake News' Impact The Stock Market?," *Forbes*, 2017. <https://www.forbes.com/sites/kenrapoza/2017/02/26/can-fake-news-impact-the-stock-market/?sh=2abae3c2fac0>.
52. S. Haider, L. Luceri, A. Deb, A. Badawy, N. Peng, and E. Ferrara, "Detecting Social Media Manipulation in Low-Resource Languages," *CoRR*, vol. abs/2011.0, 2020, [Online]. Available: <https://arxiv.org/abs/2011.05367>.
53. X. Chen, "Learning Deep Representations for Low-Resource Cross-Lingual Natural Language Processing," PhD thesis, *Cornell University*, 2019, doi: [10.7298/agpn-g679](https://arxiv.org/abs/10.7298/agpn-g679).
54. H. Q. Abonizio, J. I. de Morais, G. M. Tavares, and S. B. Junior, "Language-Independent Fake News Detection: English, Portuguese, and Spanish Mutual Features," *Futur. Internet*, vol. 12, no. 5, p. 87, 2020, doi: 10.3390/fi12050087.
55. D. Kar, M. Bhardwaj, S. Samanta, and A. P. Azad, "No Rumours Please! A Multi-Indic-Lingual Approach for COVID Fake-Tweet Detection," *CoRR*, vol. abs/2010.0, 2020, [Online]. Available: <https://arxiv.org/abs/2010.06906>.
56. X. Chen, Y. Sun, B. Athiwaratkun, C. Cardie, and K. Q. Weinberger, "Adversarial Deep Averaging Networks for Cross-Lingual Sentiment Classification," *Trans. Assoc. Comput. Linguist.*, vol. 6, pp. 557–570, 2018, [Online]. Available: <https://transacl.org/ojs/index.php/tacl/article/view/1413>.
57. J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," in *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, {NAACL-HLT} 2019, Minneapolis, MN, USA, June 2-7, 2019, Volume 1 (Long and Short Papers)*, 2019, pp. 4171–4186, doi: 10.18653/v1/n19-1423.
58. J. Libovický, R. Rosa, and A. Fraser, "How Language-Neutral is Multilingual BERT?," *CoRR*, vol. abs/1911.0, 2019, [Online]. Available: <http://arxiv.org/abs/1911.03310>.
59. Y. Wang *et al.*, "EANN: Event Adversarial Neural Networks for Multi-Modal Fake News Detection," in *Proceedings of the 24th {ACM} {SIGKDD} International Conference on Knowledge Discovery & Data Mining, {KDD} 2018, London, UK, August 19-23, 2018*, 2018, pp. 849–857, doi: 10.1145/3219819.3219903.
60. 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).
61. Narzisi, G., 2008. Yahoo! Clusty-Adding real-time clustering functionality to the Yahoo! web search engine.
62. 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).
63. 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.
 64. Ou, W. and Lin, Y., 2020. A clarifying question selection system from NTES_ALONG in Convai3 challenge. *arXiv preprint arXiv:2010.14202*.
 65. S. Zhang, L. Yao, A. Sun, and Y. Tay, “Deep learning based recommender system: A survey and new perspectives,” *ACM Comput Surv*, vol. 52, no. 1, Feb. 2019, doi: 10.1145/3285029.
 66. W. Lei, X. He, M. de Rijke, ... T. C. the 43rd I. A. S., and undefined 2020, “Conversational recommendation: Formulation, methods, and evaluation,” *dl.acm.org*, pp. 2425–2428, Jul. 2020, doi: 10.1145/3397271.3401419.
 67. Y. Deng, Y. Li, F. Sun, B. Ding, and W. Lam, “Unified Conversational Recommendation Policy Learning via Graph-based Reinforcement Learning,” *SIGIR 2021 - Proceedings of*
 68. T. Zhang, Y. Liu, P. Zhong, C. Zhang, ... H. W. preprint arXiv, and undefined 2021, “Kecrs: Towards knowledge-enriched conversational recommendation system,” *arxiv.org*, Accessed: Dec. 04, 2022. [Online]. Available: <https://arxiv.org/abs/2105.08261>
 - [69] F. Alkomah and X. Ma, “A Literature Review of Textual Hate Speech Detection Methods and Datasets,” *Information*, vol. 13, no. 6, p. 273, May 2022, doi: 10.3390/info13060273.
 - [70] F. Poletto, V. Basile, M. Sanguinetti, C. Bosco, and V. Patti, “Resources and benchmark corpora for hate speech detection: a systematic review,” *Lang. Resour. Eval.*, vol. 55, no. 2, pp. 477–523, Jun. 2021, doi: 10.1007/s10579-020-09502-8.
 - [71] S. Hinduja and J. W. Patchin, “Bullying, Cyberbullying, and Suicide,” *Arch. Suicide Res.*, vol. 14, no. 3, pp. 206–221, Jul. 2010, doi: 10.1080/13811118.2010.494133.
 - [72] U. Belavusau, “Fighting Hate Speech through EU Law,” *Amst. Law Forum*, vol. 4, no. 1, p. 20, Dec. 2012, doi: 10.37974/ALF.208.
 - [73] S. MacAvaney, H.-R. Yao, E. Yang, K. Russell, N. Goharian, and O. Frieder, “Hate speech detection: Challenges and solutions,” *PLOS ONE*, vol. 14, no. 8, p. e0221152, Aug. 2019, doi: 10.1371/journal.pone.0221152.
 - [74] W. Yin and A. Zubiaga, “Towards generalisable hate speech detection: a review on obstacles and solutions,” *PeerJ Comput. Sci.*, vol. 7, p. e598, Jun. 2021, doi: 10.7717/peerj-cs.598.
 - [75] S. Chakraborty, P. Dutta, S. Roychowdhury, and A. Mukherjee, “CRUSH: Contextually Regularized and User anchored Self-supervised Hate speech Detection.” arXiv, May 04, 2022. Accessed: Feb. 26, 2023. [Online]. Available: <http://arxiv.org/abs/2204.06389>
 - [76] P. Mishra, M. Del Tredici, H. Yannakoudakis, and E. Shutova, “Author Profiling for Hate Speech Detection.” arXiv, Feb. 14, 2019. Accessed: Feb. 26, 2023. [Online]. Available: <http://arxiv.org/abs/1902.06734>
 - [77] J. Uyheng and K. M. Carley, “Characterizing network dynamics of online hate communities

around the COVID-19 pandemic,” *Appl. Netw. Sci.*, vol. 6, no. 1, p. 20, Mar. 2021, doi: 10.1007/s41109-021-00362-x.

- [78] A. Ghosh Chowdhury, A. Didolkar, R. Sawhney, and R. R. Shah, “ARHNet - Leveraging Community Interaction for Detection of Religious Hate Speech in Arabic,” in *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics: Student Research Workshop*, Florence, Italy, 2019, pp. 273–280. doi: 10.18653/v1/P19-2038.
- [79] L. Reiners and C. Schemer, “A Feature-Based Approach to Assess Hate Speech in User Comments,” *Quest. Commun.*, no. 38, pp. 529–548, Dec. 2020, doi: 10.4000/questionsdecommunication.24808.