

تولید توضیح شخصی سازی شده برای سیستم پیشنهاددهنده لیست توئیت مبتنی بر شباهت معنایی هشتگ‌ها

حوا علیزاده نوقابی، بهشید بهکمال^{*}، صالحه ناصری، محسن کاهانی
گروه مهندسی کامپیوتر، دانشگاه فردوسی مشهد
behkamal@um.ac.ir

چکیده

امروزه سیستم‌های پیشنهاددهنده‌ی لیست‌های توئیت با بکارگیری اطلاعات مختلف کاربران و لیست‌ها و همچنین اعمال الگوریتم‌های پیچیده، توانسته‌اند به دقت بالایی در پیش‌بینی برسند و پیشنهادهای مرتبط برای هر کاربر تولید کنند، اما قابلیت توضیح‌پذیری در این سیستم‌ها به عنوان یک چالش مطرح می‌باشد. یک توضیح مناسب به همراه لیست توئیت پیشنهادشده علاوه بر افزایش اعتماد و رضایت کاربران، می‌تواند به آن‌ها در تصمیم‌گیری آگاهانه کمک نماید. از این رو در این مقاله، یک مدل تولید توضیح، ارائه می‌شود که به صورت خودکار، یک توصیف شخصی سازی شده برای لیست توئیت پیشنهادشده ایجاد می‌نماید. بطور دقیق‌تر، این مدل با انتخاب مجموعه‌ای از هشتگ‌های پرتکرار لیست توئیت که شباهت معنایی با تاریخچه فعالیت‌های قبلی کاربر دارد، سعی می‌کند موضوع لیست را به گونه‌ای قابل درک در قالب یک توضیح همراه با لیست پیشنهادشده به کاربر نمایش دهد. پس از جمع‌آوری یک مجموعه داده واقعی از شبکه اجتماعی توئیت، با انجام آزمایش‌ها نشان داده شد که مدل پیشنهادی قادر به تولید توضیح برای درصد بالایی از پیشنهادهای ایجادشده براساس یک مدل پیشنهاددهنده‌ی پایه می‌باشد.

واژگان کلیدی: سیستم‌های پیشنهاددهنده توضیح‌پذیر، توضیح شخصی سازی شده، لیست توئیت، شباهت معنایی هشتگ‌ها

۱. مقدمه

مطالب آن در زمینه‌ی مشخص شده استفاده کنند. در همین راستا سیستم‌های پیشنهاددهنده‌ی لیست توئیت [3]، [2] طراحی شدند که با استفاده از ویژگی‌های مختلف کاربر و همچنین ویژگی‌های لیست، سعی در بهبود عملکرد در جهت ارائه پیشنهادات مرتبط برای هر کاربر دارند و اما چالش اصلی آن‌ها عدم فراهم نمودن توضیح^۳ مناسب برای هر پیشنهاد می‌باشد، عبارتی قابلیت توضیح‌پذیری در این سیستم‌ها اهمیت بسیاری دارد که به عنوان یک چالش مهم مطرح است.

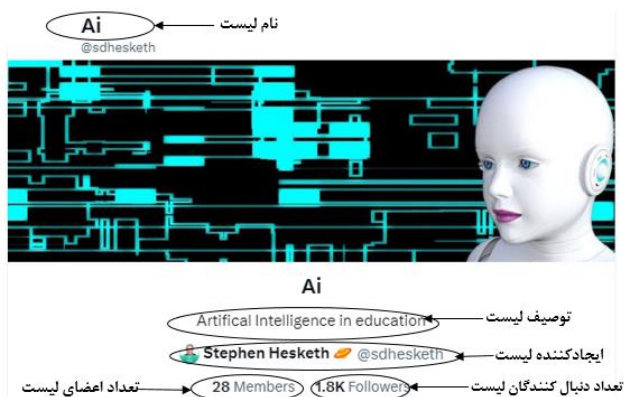
در سال‌های اخیر فاکتورهایی مانند اعتماد و رضایت کاربران در شبکه‌های اجتماعی بسیار مورد توجه قرار گرفته است و تحت‌تأثیر عواملی مانند قابلیت توضیح‌پذیری در سیستم‌های پیشنهاددهنده

افزایش قابل توجه محتوای تولید شده در شبکه‌های اجتماعی، باعث شده است که جهت سازماندهی اطلاعات مرتبط در یک زمینه، راه‌حلهایی توسط پلتفرم‌های مختلف به کاربران ارائه داده شود. به عنوان مثال، لیست‌ها در توئیت^۱ برای مقابله با مشکل سرشار اطلاعات در سال ۲۰۱۳ معرفی شد [1]. یک لیست توئیت^۲، گروهی از حساب‌های کاربری است که به موضوع آن لیست علاقه‌مند هستند. همانطور که در شکل ۱ نشان داده می‌شود هر لیست تعدادی عضو دارد که امکان درج توئیت در مورد موضوع آن لیست را دارند و سایر کاربران می‌توانند لیست را دنبال کنند و از

^۱ اخیراً این شبکه اجتماعی، نام جدید X گرفته است.

^۲ Twitter List

^۳ explanation



شکل ۱ - یک نمونه لیست توئیتر

به صورت مشخص‌تر می‌توان گفت از آنجا که هر کاربر در توئیتر دارای یک تاریخچه‌ی پست می‌باشد، اگر توضیح ارائه شده، بجای بکارگیری تمام هشتگ‌های پرتکرار لیست پیشنهادشده، تنها موارد مرتبط با تاریخچه‌ی پست‌های کاربر را دربرگیرد، دو مزیت به همراه خواهد داشت: (۱) توضیح برای کاربر، قابل فهم^۶ خواهد شد که فراهم نمودن توضیح قابل فهم در **بافتار**^۷ سیستم‌های پیشنهاددهنده توضیح‌پذیر اهمیت دارد [10]، به عبارتی هر توضیح براساس میزان دانش هر کاربر می‌تواند شخصی‌سازی شود تا کاربر قادر باشد شناخت بهتری از قلم پیشنهادشده با کمک توضیح فراهم شده به دست آورد [11]. (۲) این چنین توضیحی می‌تواند به متقاعدسازی کاربر نسبت به چرایی پیشنهاد این لیست توسط سیستم پیشنهاددهنده، کمک کند. از آنجا که این توضیح به ارتباط لیست پیشنهادشده با فعالیت‌های قبلی کاربر (در اینجا تعاملات قبلی از طریق پست‌های توئیتر) می‌پردازد، کاربر دلیل پیشنهاد چنین قلمی از سمت سیستم پیشنهاددهنده را درک می‌کند و به مدل پیشنهاددهنده اعتماد می‌کند که یکی از اهداف مهم مدل‌های تولید توضیح است [6].

در راستای شخصی‌سازی توضیح، برای محاسبه شباهت بین هشتگ‌های لیست توئیتر پیشنهادشده و هشتگ‌های استخراج شده از تاریخچه پست‌های کاربر از روش مبتنی بر BERT [12]، استفاده می‌شود که قادر است بصورت معنایی، ارتباط بین هشتگ‌ها را مشخص نماید. بنابراین اهداف اصلی این مقاله در سه مورد زیر خلاصه می‌شود:

- ارائه یک روش تولید توضیح توصیفی برای لیست‌های توئیتر پیشنهادشده به کاربر، به گونه‌ای که به تصمیم‌گیری آگاهانه کمک نماید.
- شخصی‌سازی توضیح توصیفی برای افزایش قابلیت فهم آن که براساس شباهت معنایی بین هشتگ‌های لیست

اجتماعی می‌باشد [4]. در واقع ارائه یک توضیح مناسب همراه هر قلم پیشنهادی (مانند یک لیست توئیتر)، علاوه بر جلب رضایت کاربران، باعث تعامل بیشتر آن‌ها با پیشنهاد ارائه شده می‌شود. از این رو هدف اصلی این تحقیق فراهم نمودن توضیحی قابل فهم و شخصی‌سازی شده می‌باشد. از آنجا که توسعه روش‌های تولید توضیح تعقیبی^۱ [7]، [6]، [5]، که در آن‌ها توضیح‌ها پس از ایجاد پیشنهادها و بصورت مستقل از مدل پیشنهاددهنده، تولید می‌شود، اخیراً مورد توجه زیادی قرار گرفته است، این تحقیق مدل تولید توضیحی را ارائه می‌دهد که دارای انعطاف‌پذیری برای بکارگیری طیف وسیعی از مدل‌های پیشنهاددهنده می‌باشد. به عبارتی هر نوع مدل پیشنهاددهنده‌ی لیست توئیتر می‌تواند از مدل تولید توضیح ارائه شده در این مقاله برای فراهم نمودن توضیح استفاده نماید.

توضیح‌ها در یک سیستم پیشنهاددهنده ممکن است برای رسیدن به اهداف مختلفی مانند شفافیت^۲، متقاعدسازی^۳، اثربخشی^۴، رضایت‌مندی^۵ و غیره فراهم شده باشند [9]، [8]. به عنوان مثال یک توضیح ممکن است دلیل پیشنهاد یک قلم خاص به کاربر را بیان کند: "این قلم [لنز] به شما پیشنهاد شده است چون شما قبلاً قلم مرتبطی [دوربین] خریداری کرده‌اید." [6]. بطور مشابه هدف یک توضیح می‌تواند فراهم نمودن اطلاعات در مورد قلم پیشنهادشده برای کمک به اتخاذ تصمیم آگاهانه توسط کاربر باشد. به عنوان مثال بیان یک توضیح حاوی اطلاعات راجع به محتوا و موضوع یک لیست توئیتر می‌تواند به کاربر در شناخت آن لیست و در نتیجه تصمیم‌گیری برای پذیرش یا رد آن پیشنهاد کمک بسزایی نماید. با توجه به اینکه در اکثر مواقع حدس زدن موضوع یا موضوعات اصلی یک لیست توئیتر فقط با دیدن نام آن برای کاربران سخت است، معمولاً کاربران به صورت دستی تعدادی از توییت‌های لیست را بررسی می‌کنند تا موضوع آن را درک کنند که خسته‌کننده و زمان‌بر است. بنابراین تولید خودکار یک توصیف برای یک لیست توئیتر می‌تواند یک نوع توضیح آگاهی‌دهنده باشد.

از آنجا که هشتگ‌ها در شبکه‌های اجتماعی به عنوان کلمات کلیدی بسیار استفاده می‌شوند و به موضوع یا مفهوم خاصی در یک پست (توئیت) اشاره می‌کنند، در این تحقیق توصیف یک لیست توئیتر براساس هشتگ‌های پرتکرار محتوای آن در نظر گرفته می‌شود. اگرچه که استخراج هشتگ‌های پرتکرار لیست و ارائه آن به کاربر می‌تواند توصیفی مناسب راجع به لیست فراهم نماید، اما برآنیم تا این توصیف، برای کاربر مدنظر، شخصی‌سازی شود.

^۱ post-hoc explanation

^۲ transparency

^۳ persuasiveness

^۴ effectiveness

^۵ satisfaction

^۶ understandable

^۷ context

توئیر پیشنهاددهنده و هشتگ‌های موجود در تاریخچه پست‌های کاربر انجام می‌شود.

- یک مجموعه داده واقعی از توئیر جمع‌آوری می‌شود و جهت ارزیابی روش پیشنهادی، آزمایش‌های مختلف انجام می‌شود.

در ادامه مقاله، در بخش ۲ به بررسی کارهای گذشته‌ی مرتبط با موضوع تحقیق پرداخته می‌شود و سپس روش پیشنهادی به همراه جزئیات آن در بخش ۳ شرح داده می‌شود. نتایج آزمایش‌ها و تحلیل آن‌ها در بخش ۴ آورده شده و در بخش ۵، جمع‌بندی آمده است.

۲. کارهای گذشته

امروزه توضیح‌پذیر بودن مدل‌های پیشنهاددهنده یک مساله حائز اهمیت به شمار می‌رود. عبارت دیگر یک سیستم پیشنهاددهنده علاوه بر آن که عملکرد خوبی در تولید پیشنهادهای مرتبط به کاربران دارد، بایستی توضیحات مفید به همراه اقلام پیشنهاددهنده به کاربران ارائه نماید. مدل‌های تولید توضیح را می‌توان به دو دسته کلی طبقه‌بندی کرد: مستقل از مدل پیشنهاددهنده و یا وابسته به مدل پیشنهاددهنده [13], [8], [6]. رویکردهای دسته اول با استفاده از مدل‌های جداگانه‌ای که مستقل از الگوریتم پیشنهاددهنده هستند، توضیحاتی را ایجاد می‌کنند که به آن‌ها توضیح تعقیبی می‌گویند. اینگونه رویکردها اغلب انعطاف‌پذیر هستند و می‌توانند برای طیف وسیعی از الگوریتم‌های پیشنهاددهنده اعمال شوند. تکنیک‌های متداول برای ایجاد توضیحات تعقیبی شامل مدل‌های تقلیدکننده^۱ [14]، مدل‌های مبتنی بر روش‌های داده‌کاوی مثل قانون‌کاوی انجمنی^۲ [15] و یا مبتنی بر زیرگراف [16] است. از طرف دیگر توضیحات ذاتی^۳ توسط مدل پیشنهاددهنده ایجاد می‌شوند و اغلب هدف آن‌ها ارائه توضیحی در مورد اینکه چرا الگوریتم پیشنهاددهنده یک قلم خاص را به کاربر مدنظر پیشنهاد داده است، می‌باشد. از جمله تحقیقات انجام شده در این زمینه می‌توان روش‌های مبتنی بر مدل‌های فاکتورسازی [18], [17]، مدل‌های مبتنی بر گراف دانش [19], [22], [21], [20]، مدل‌های یادگیری عمیق [25], [24], [23] و مدل‌های مبتنی بر قاعده [27], [26] را نام برد.

همچنین توضیح‌پذیر بودن سیستم‌های پیشنهاددهنده اجتماعی در ایجاد اعتماد کاربران بسیار موردتوجه قرار گرفته است، زیرا برای حفظ پایداری شبکه‌های اجتماعی، ضروری دانسته شده است [28]. در برخی از تحقیقات قبلی مرتبط با تولید توضیح برای

پیشنهادهای اجتماعی، توضیحات را بر اساس شبکه اجتماعی کاربر ارائه می‌دهند. به عنوان مثال، وانگ و همکاران [29] توضیحی را فراهم می‌کنند که از قالب "کاربر آ و کاربر ب نیز این قلم را می‌پسندند." پیروی می‌کند. برای ایجاد این توضیح، آن‌ها الگوریتمی را برای شناسایی مجموعه بهینه کاربران برای گنجاندن در توضیح طراحی کردند که کاربر مدنظر را در مورد قلم پیشنهاددهنده متقاعد کنند. در برخی دیگر از مدل‌های توضیح‌پذیر، یک شبکه اطلاعاتی ناهمگون با توجه به اطلاعات کاربران و اقلام ایجاد می‌شود و سپس توضیحات به شکل مسیر^۴ اتصال کاربر به قلم پیشنهاددهنده در گراف ساخته شده شامل تعاملات قبلی کاربران با اقلام و همچنین ویژگی‌های اقلام ارائه می‌شود [31], [30]. اخیراً، محققان محتوای تولید شده توسط کاربران مانند نظرات^۵ را در تولید توضیح بکار می‌گیرند [33], [32]. به عنوان مثال، رن و همکاران [33] مدلی را معرفی کرده‌اند که از دیدگاه‌ها^۶ به عنوان توضیح استفاده می‌کند. این دیدگاه‌ها با ترکیبی از مفهوم^۷، موضوع^۸ و برچسب احساس^۹ نشان داده می‌شوند که هم از نظرات کاربران و هم از ارتباطات اجتماعی آن‌ها استخراج می‌شود.

در مقایسه با مدل‌های متنوع تولید توضیح برای اقلام پیشنهاددهنده، روش این مقاله بر ارائه توضیحی تعقیبی و آگاهی‌دهنده برای لیست توئیر تأکید دارد. عبارت دیگر، مدل ارائه شده می‌تواند مستقل از مدل پیشنهاددهنده و پیچیدگی‌های ذاتی الگوریتم آن، به ایجاد خودکار توضیحات شخصی‌سازی شده مبتنی بر هشتگ بپردازد. چنین توضیحاتی به کاربر در شناخت لیست توئیر پیشنهاددهنده، کمک خواهد کرد و نقش قابل توجهی در تصمیم‌گیری کاربر در خصوص رد یا پذیرش لیست پیشنهاددهنده خواهد داشت.

۳- تولید توضیح شخصی‌سازی شده برای لیست‌های توئیر

اگر U مجموعه کاربران و L مجموعه لیست‌های توئیر باشد، آن گاه برای هر لیست $l \in L$ با توجه به محتوای آن می‌توان مجموعه هشتگ‌های پرتکرار را استخراج کرده و به شکل $T_l = \{t_l^1, t_l^2, \dots, t_l^m\}$ در نظر گرفت. از طرف دیگر برای هر کاربر $u \in U$ با توجه به تاریخچه‌ی پست‌های او می‌توان هشتگ‌های پرتکرار را به شکل $T_u = \{t_u^1, t_u^2, \dots, t_u^n\}$ نمایش داد. آنگاه می‌توان مساله این تحقیق که هدف آن تولید توضیح شخصی‌سازی شده برای لیست‌های پیشنهاددهنده است را به شکل زیر تعریف نمود.

^۴ path

^۵ reviews

^۶ viewpoints

^۷ concept

^۸ topic

^۹ sentiment label

^۱ surrogate models

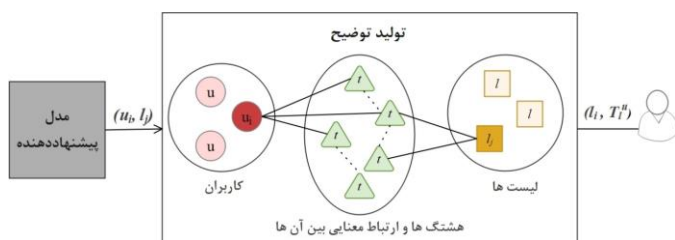
^۲ association rule mining

^۳ intrinsic

۰.۷ تنظیم گردید. پس از تولید توضیح، می‌توان آن را در قالب‌های مختلف به کاربر نمایش داد که در این تحقیق همانند [34], [35] از ابرکلمه^۱ برای نمایش هشتگ‌ها استفاده می‌شود.

۴- آزمایش‌ها و ارزیابی نتایج

۴-۱- مجموعه داده. برای انجام آزمایش‌ها، یک مجموعه داده از توئیتر شامل کاربران و لیست‌هایی که هر کاربر دنبال می‌کند، با استفاده از Tweepy^۲ جمع‌آوری گردید. مشابه [36]، از Ashton^۳، Kutcher^۴، بازیگر و کارآفرین محبوب توئیتر با بیش از ۱۷ میلیون دنبال‌کننده، شروع شد و سپس با توجه به لیست‌های او، برای افزایش کاربران مجموعه داده، از هر لیست حداکثر ۵۰ دنبال‌کننده بصورت تصادفی انتخاب شدند. بعد از آن با توجه به کاربران جدید، مجموعه لیست‌ها با جمع‌آوری حداکثر ۱۰ لیست به ازای هر کاربر جدید، گسترش داده شد. این چرخه‌ی گسترش کاربران و لیست‌ها ۴ بار تکرار شد که منجر به مجموعه داده‌ای با ۲۵,۷۸۸ کاربر و ۱۷,۶۰۴ لیست و ۶۶,۵۸۳ رابطه‌ی کاربر-لیست شد.



شکل ۲- روش تولید توضیح مبتنی بر هشتگ برای لیست توئیتر

Algorithm: Tag-based Explanation Generation
Input: (user u , recommended List l), frequent tags of u (T_u), frequent tags of l (T_l)
Output: the subset of l 's tags ($T_l^u \subset T_l$)

```

 $T_l^u = \emptyset$ 
for  $t_l$  in  $T_l$  do
    for  $t_u$  in  $T_u$  do
         $\text{sim} = \text{BERT.get\_semantic\_similarity}(t_l, t_u)$ 
        if  $\text{sim} > \text{threshold}$  and  $t_l$  not in  $T_l^u$  then:
             $T_l^u.append(t_l)$ 
        end if
    end for
end for
return  $T_l^u$ 

```

شکل ۳- الگوریتم تولید توضیح مبتنی بر هشتگ

تعریف مساله. فرض کنید به کاربر $u \in U$ ، یک لیست توئیتر $l \in L$ ، توسط یک مدل پیشنهاددهنده، پیشنهاد شده است، آنگاه یافتن زیرمجموعه‌ی مناسب از هشتگ‌های لیست l ، به گونه‌ای که برای کاربر u با در نظر داشتن مجموعه هشتگ‌های او یعنی T_u ، شخصی‌سازی شده باشد، مد نظر است. این زیرمجموعه از هشتگ‌ها یعنی $T_l^u \subset T_l$ به عنوان یک توضیح توصیفی برای لیست پیشنهاد شده‌ی l در نظر گرفته می‌شود.

شکل ۲ بخش‌های مدل تولید توضیح برای لیست‌های توئیتر پیشنهادشده به کاربران را نشان می‌دهد. از آنجایی که در این روش، توضیحات از نوع تعقیبی هستند، مدل پیشنهاددهنده به شکل جعبه سیاه در شکل ۲ رسم شده است، زیرا برای هر لیست پیشنهادشده بدون در نظر گرفتن الگوریتم پیشنهاددهنده، یک توضیح فراهم خواهد شد. بنابراین ورودی کامپوننت تولید توضیح شامل جفت (کاربر u ، لیست پیشنهادشده l) می‌باشد. آنگاه براساس این ورودی و همچنین هشتگ‌های پرتکرار کاربر و لیست و ارتباط معنایی بین آن‌ها، یک توضیح تولید می‌شود که همراه با لیست پیشنهادشده، یعنی جفت (l, T_l^u) ، به کاربر ارائه داده شود و در نتیجه قابلیت توضیح‌پذیری به مدل پیشنهاددهنده افزوده شود.

همانطور که در کامپوننت تولید توضیح در شکل ۲ مشاهده می‌شود نیاز به دانستن ارتباط معنایی بین هشتگ‌ها (خط چین) می‌باشد. از آنجایی که مدل BERT، یک مدل مطرح در زمینه پردازش زبان طبیعی است و به عنوان یکی از پیشروهای این حوزه به حساب می‌آید، برای استخراج روابط معنایی بین هشتگ‌ها از آن استفاده می‌شود. در واقع این مدل می‌تواند مفاهیم پیچیده و ارتباطات معنایی کلمات را با دقت خوبی مشخص کند و بنابراین در این مقاله به عنوان یک ابزار قدرتمند مورد استفاده قرار می‌گیرد تا شباهت معنایی بین دو هشتگ را تعیین کند.

همانطور که در تعریف مساله بیان شد انتخاب زیرمجموعه‌ی مناسبی از هشتگ‌های پرتکرار لیست با در نظر داشتن هشتگ‌های موردعلاقه و شناخته شده برای کاربر، منجر به تولید توضیح شخصی‌سازی شده می‌شود. الگوریتم آورده شده در شکل ۳ نحوه انتخاب این زیرمجموعه را بیان می‌کند. بطور دقیق‌تر می‌توان گفت اگر شباهت معنایی بین یک هشتگ پرتکرار از لیست و یک هشتگ پرتکرار کاربر از یک مقدار آستانه (threshold) بیشتر باشد به عنوان عضوی از مجموعه هشتگ‌های خروجی الگوریتم خواهد بود. مقدار این آستانه می‌تواند به هر عددی بین صفر و ۱ تنظیم شود که هر چه این عدد به ۱ نزدیکتر باشد، یعنی شباهت هشتگ‌های خروجی الگوریتم از نظر معنایی بسیار شبیه به هشتگ‌های قبلی کاربر خواهد بود. در آزمایشات این مقاله با توجه به اعداد مشاهده شده مربوط به شباهت هشتگ‌ها براساس BERT، مقدار آستانه به

^۱ word-cloud

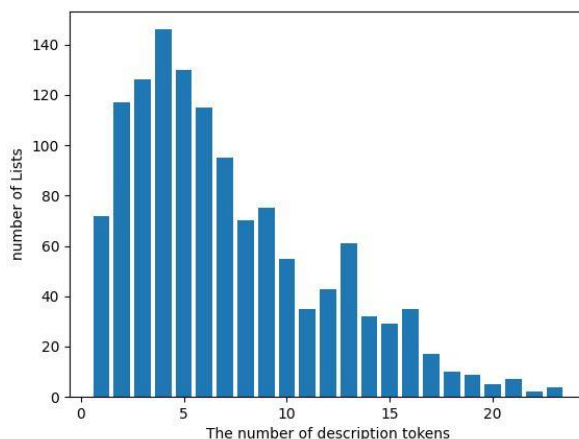
^۲ <https://www.tweepy.org/>

شود و سپس بر اساس شباهت بین پروفایل‌ها برای هر کاربر، k لیست مرتبط به او پیشنهاد داده شود. ساخت پروفایل‌ها مبتنی بر روش‌های مدل‌سازی موضوع^۱، انجام می‌شود، بدین طریق که پس از اعمال یک روش مدل‌سازی موضوع و استخراج Z موضوع نهفته در متون (در این تحقیق توئیت‌ها)، برای هر کاربر و هر لیست، موضوع‌ها و وزن هر کدام برای هر کاربر و هر لیست مشخص می‌شود. بر این اساس دو ماتریس ساخته می‌شود: ماتریس کاربر-موضوع به شکل $A = |U| \times |Z|$ و ماتریس موضوع-لیست یعنی $B = |Z| \times |L|$. برای تولید پیشنهادات، ماتریس جدیدی مبتنی بر این دو ماتریس به صورت $R = A \times B$ ایجاد می‌شود و با توجه به آن میتوان k -برترین پیشنهادات را برای هر کاربر مشخص کرد.

برای اعمال الگوریتم پیشنهاددهنده روی مجموعه داده این تحقیق، یک روش مدل‌سازی موضوع جدید به نام BERTopic^۲ بکار گرفته شد. روش BERTopic ابتدا با استفاده از مدل‌های زبانی مبتنی بر ترانسفورماتور از قبل آموزش‌دیده، تعبیه^۳ متون را بدست می‌آورد. در گام بعدی، پس از کاهش ابعاد تعبیه، آنها را خوشه‌بندی^۳ می‌کند و در نهایت موضوعات و کلمات متناظر آن را با روش C-TF-IDF ارائه می‌دهد [39], [38]. در جدول ۲ تعدادی از موضوعات و کلمه‌های متناظر آن که توسط BERTopic از مجموعه داده این تحقیق استخراج شده است، نشان داده می‌شود.

جدول ۲- برخی از موضوعات استخراج شده از مجموعه داده

موضوع	کلمات موضوع
AI	'ai', 'machinelearning', 'artificialintelligence', 'datascience', 'robot', 'bigdata', 'ml', 'learning', 'machine learning', 'artificial'
Music	'music', 'song', 'album', 'kanye', 'songs', 'wizkid', 'spotify', 'tickets', 'dj', 'dance'
Cybersecurity	'cybersecurity', 'cyber', 'security', 'infosec', 'cyber security', 'hacking', 'cyberwarrior', 'cybercrime', 'privacy', 'malware'
Cancer	'cancer', 'breast', 'breast cancer', 'heart', 'prostate', 'patients', 'prostate cancer', 'treatment', 'disease', 'surgery'
Bitcoin	'bitcoin', 'crypto', 'btc', 'blockchain', 'defi', 'ethereum', 'cryptocurrency', 'dip', 'market', 'buy'



شکل ۴- توزیع لیست‌های توئیت براساس تعداد توکن‌های توصیف

جدول ۱- اطلاعات آماری در مورد مجموعه داده

#کاربران	# لیست			#رابطه کاربر-لیست
۱۲۸۴	به ازای یک کاربر			۶۰۵۹
	مجموع			
	متوسط	حداقل	حداکثر	
	۵	۱	۱۰	۱۵۳۵

سپس، توئیت‌های کاربران و لیست‌ها نیز جمع‌آوری گردید، به گونه‌ای که برای هر کاربر حداکثر ۲۰۰۰ توئیت و برای هر لیست حداکثر ۵۰۰ توئیت اخیر گردآوری شدند. مشابه [37]، برای تمرکز بر کاربران فعال توئیت، فقط کاربرانی که حداقل ۲۰۰ توئیت داشته‌اند، به عنوان کاربران نهایی در نظر گرفته شد. پس از این پیش‌پردازش‌ها، مجموعه داده شامل ۱۲۸۴ کاربر، ۱۵۳۵ لیست و ۶۰۵۹ رابطه کاربر-لیست می‌باشد. جدول ۱ برخی از آمارهای اصلی در مورد مجموعه داده را نمایش می‌دهد. علاوه بر آن با بررسی زیرمجموعه تصادفی شامل ۳۰۰۰ لیست، مشخص می‌شود که برای ۶۱٪ آنها توصیفی توسط ایجادکننده لیست درج نشده است (مشابه توصیف ذکر شده در مثال شکل ۱) و برای ۳۹٪ درصد باقیمانده، توزیع لیست‌های توئیت بر اساس تعداد توکن‌های توصیف آن در شکل ۴ مشاهده می‌شود که بیان می‌کند که طول توصیف لیست در بیشتر موارد کوتاه است. بنابراین، ارائه توضیحات با مقدار اطلاعات کافی برای کمک به کاربران برای شناخت بهتر لیست ضروری می‌باشد.

۲-۴- مدل پیشنهاددهنده. همانطور که در بخش ۳ گفته شد تمرکز این تحقیق بر ارائه توضیحات تعقیبی برای لیست‌هایی است که توسط یک مدل پیشنهاددهنده تولید شده است. از این رو برای تولید پیشنهادات برای هر کاربر، یک مدل پیشنهاددهنده لیست توئیت معرفی شده در [2] بکار گرفته می‌شود. در این مدل، ابتدا باید پروفایلی برای کاربران و لیست‌ها براساس محتوای آنها ساخته

^۱ topic modeling
^۲ embedding
^۳ clustering

جدول ۳- چند مثال از شباهت معنایی بین هشتگ‌ها

شباهت معنایی	جفت هشتگ
0.7101	(#ArtificialIntelligence, #Python)
0.8873	(#deepLearning, #machineLearning)
0.8190	(#AI, #ML)
0.5544	(#football, #deepLearning)
0.5946	(#politics, #algorithm)

۳-۴- شباهت معنایی بین هشتگ‌ها. همانطور که در بخش ۳ گفته شد برای تعیین توضیح مبتنی بر هشتگ برای یک لیست توئیت، نیاز است تا ارتباط معنایی بین هشتگ‌های لیست و هشتگ‌های کاربری که این لیست به او پیشنهاد شده است، مشخص گردد. برای شباهت معنایی بین دو هشتگ، BERT بکار گرفته می‌شود که در جدول ۳، چند زوج هشتگ و میزان شباهت آن‌ها آورده شده که بیانگر عملکرد قابل قبول BERT می‌باشد.

۴-۴- ارزیابی مدل تولید توضیح

در این تحقیق یک مدل جهت تولید توضیح مبتنی بر هشتگ برای لیست‌های توئیت ارائه شد. مهمترین معیار برای ارزیابی یک مدل تولید توضیح تعقیبی، معیار وفاداری^۱ [15] می‌باشد، که بیانگر درصد اقلام پیشنهادشده‌ای است که مدل می‌تواند برای آن‌ها توضیح ایجاد کند. عبارت دیگر یک مدل تولید توضیح باید قادر باشد برای درصد قابل توجهی از پیشنهادات، توضیح فراهم نماید.

معادله ۱ نحوه محاسبه معیار وفاداری را بیان می‌کند، که در آن R_u اقلام پیشنهادشده به کاربر u با طول k است. برای هر کاربر، تعداد اقلام پیشنهادشده‌ای که توضیح متناظر با آن، یعنی e_i ، مخالف تهی باشد، بر تعداد کل اقلام (k) تقسیم می‌شود. آنگاه میانگین مقادیر برای تمام کاربران به عنوان مقدار این معیار خواهد بود.

$$fidelity = \frac{1}{|U|} \sum_{u \in U} \frac{|\{r_i \in R_u | e_i \neq \emptyset\}|}{K} \quad \text{معادله (۱)}$$

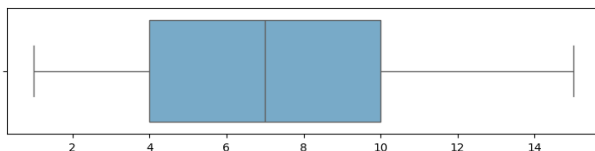
در جدول ۴ به ازای k های مختلف که در آن k تعداد پیشنهاد ارائه شده به هر کاربر توسط مدل پیشنهاددهنده‌ی پایه (بخش ۴-۲) می‌باشد، مقدار معیار وفاداری آورده شده است. به عنوان مثال در حالت $k=1$ که به هر کاربر فقط یک قلم (در اینجا یک لیست توئیت) توصیه می‌شود، اگر مدل تولید توضیح نتواند توضیح شخصی‌سازی شده‌ی مبتنی بر هشتگ برای این قلم فراهم کند، متریک وفاداری برای این کاربر برابر با صفر خواهد بود. در واقع جدول ۴، میانگین مقدارهای وفاداری مربوط به تمام کاربران

مجموعه داده را به عنوان مقدار نهایی وفاداری برای هر k بیان می‌کند. همانطور که نتایج نشان می‌دهد، مدل تولید توضیح ارائه شده می‌تواند به ازای k های مختلف، بطور میانگین برای ۹۰٫۶۲ درصد لیست‌های توئیت پیشنهادشده، توضیح فراهم نماید. در نتیجه می‌توان گفت که مدل ارائه شده در این مقاله قادر است برای درصد بالایی از لیست‌های توئیت پیشنهادشده، یک توضیح دارای هشتگ‌های پرتکرار لیست ارائه دهد به گونه‌ای که ارتباط معنایی با تاریخچه پست‌های کاربر نیز دارد. علاوه بر این، نمودار جعبه‌ای شکل ۵ براساس تعداد هشتگ‌های هر توضیح رسم شده که بیانگر آن است که میانگین تعداد هشتگ‌های هر توضیح تقریباً ۷ است.

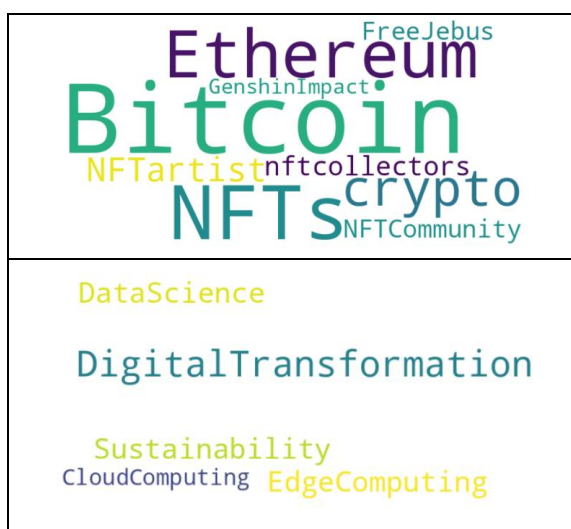
در شکل ۶ توضیح‌ها برای دو لیست توئیت با نام‌های 'crypto' و 'technology' در قالب ابرکلمه آورده شده است که کاربر می‌تواند با نگاهی به این ابرکلمات در مورد محتوا و موضوعات مطرح شده در محتوای لیست پیشنهادشده، اطلاعاتی بدست آورد که مرتبط با فعالیت‌های قبلی خود کاربر نیز می‌باشد. لازم به ذکر است که این ابرکلمه‌ها از نوع مبتنی بر فراوانی^۲ می‌باشند که در آن اندازه هر کلمه، متناسب با تعداد تکرار آن است که برای این تحقیق، برابر با فراوانی هشتگ در لیست توئیت پیشنهادشده می‌باشد.

جدول ۴- ارزیابی براساس معیار وفاداری

	k=5	k=3	k=1	
معیار وفاداری	90.75 %	90.86 %	90.24 %	



شکل ۵- طول هر توضیح براساس تعداد هشتگ



شکل ۶- ارائه توضیح به کاربر در قالب ابرکلمه

^۲ Frequency-based word-cloud

^۱ fidelity

۵-۴- مطالعه موردی

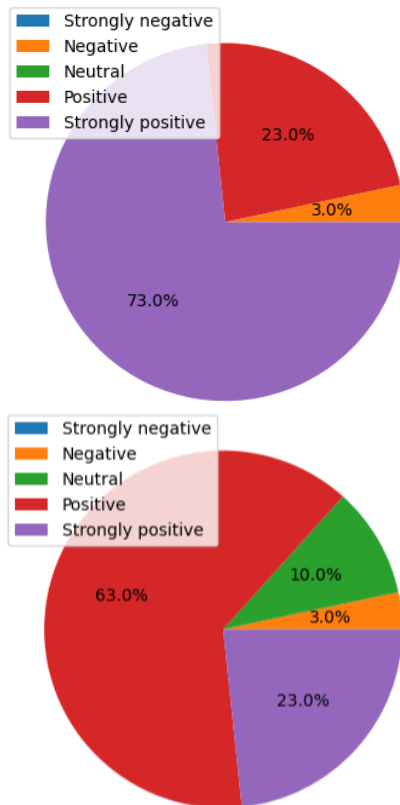
برای ارزیابی کیفی توضیحات، همانند [18] با طرح پرسش‌های مناسب بررسی می‌شود که آیا مدل تولید توضیح، موفق شده است در راستای اهداف خود، توضیحات مفیدی فراهم نماید. بر همین اساس در این تحقیق دو پرسش اصلی در نظر گرفته شده است: (۱) آیا توضیح فراهم شده باعث شناخت بیشتر لیست توئیت‌ر پیشنهادشده می‌شود؟ (۲) آیا کاربر با مشاهده توضیح، می‌تواند بینشی راجع به اینکه چرا این لیست توئیت‌ر به او پیشنهاد شده است، بدست آورد؟ در واقع اولین پرسش به ارزیابی جنبه‌ی آگاهی‌بخش بودن توضیحات اشاره دارد و مورد دوم بررسی می‌کند که آیا توضیحات شخصی‌سازی‌شده می‌تواند باعث افزایش اعتماد کاربران به سیستم پیشنهاددهنده شود، زیرا که هر توضیح، ارتباط بین قلم پیشنهادشده و فعالیت‌های قبلی کاربر را نیز در نظر دارد.

برای انجام این مطالعه موردی، به طور تصادفی ۱۰ کاربر از مجموعه داده جمع‌آوری‌شده انتخاب گردید و با کمک مدل پیشنهاددهنده پایه برای هر کدام ۳ لیست توئیت‌ر به عنوان اقلام پیشنهادی فراهم شد. سپس مدل تولید توضیح برای این اقلام، توضیحات مبتنی بر هشتگ را ایجاد نمود. آنگاه ۲ ارزیاب انسان مسلط به زبان انگلیسی، با در دست داشتن اطلاعات مربوط به کاربران و لیست‌های توئیت‌ر، به ارزیابی توضیحات براساس دو پرسش فوق پرداختند. قابل بیان است که برای هر پرسش برگرفته از [18]، ۵ گزینه پاسخ شامل به شدت منفی، منفی، خنثی، مثبت و به شدت مثبت وجود دارد.

براساس نتایج ارزیابی که در شکل ۷ مشاهده می‌شود، با توجه به پاسخ‌های مربوط به پرسش ۱ (بخش بالای شکل ۷) می‌توان گفت که توضیح تولید شده توسط مدل پیشنهادی در اکثر مواقع توانسته است لیست توئیت‌ر را به خوبی توصیف کند تا کاربران شناخت مناسبی از قلم پیشنهادشده پیدا کنند. همچنین با توجه به پاسخ‌های مربوط به پرسش ۲ (بخش پایین شکل ۷) در درصد قابل توجهی از مواقع، ارتباط بین توضیح فراهم شده با فعالیت‌های قبلی کاربران، باعث ایجاد بینشی راجع به چرایی پیشنهاد این لیست توئیت‌ر شده است.

۵- جمع‌بندی

با توجه به اهمیت قابلیت توضیح‌پذیری در سیستم‌های پیشنهاددهنده‌ی امروزی، در این تحقیق یک مدل تولید توضیح برای لیست‌های توئیت‌ر پیشنهادشده به کاربران شبکه اجتماعی ارائه شد. با توجه به اینکه هدف از فراهم نمودن توضیح در این تحقیق، ارائه اطلاعات بیشتر در مورد لیست توئیت‌ر پیشنهادشده جهت تصمیم‌گیری آگاهانه کاربران است، یک توصیف شخصی‌سازی شده مبتنی بر هشتگ در قالب ابرکلمه ایجاد می‌شود.



شکل ۷- نتایج مطالعه موردی (پاسخ‌های مربوط به پرسش اول و دوم به ترتیب در بخش بالا و پایین شکل آمده است).

برای انجام آزمایش‌ها و بررسی قابلیت عملیاتی بودن مدل تولید توضیح ارائه شده، یک مجموعه داده واقعی از توئیت‌ر جمع‌آوری گردید و سپس نتایج نشان دادند که مدل پیشنهادشده می‌تواند برای درصد قابل توجهی از لیست‌های توئیت‌ر، توضیح مختص کاربر مدنظر فراهم نماید. علاوه بر این طبق مطالعه موردی انجام شده، توضیحات فراهم شده علاوه بر اینکه منجر به شناخت بهتر لیست توئیت‌های پیشنهادشده می‌شوند، منجر به افزایش اعتماد کاربران از طریق ایجاد بینشی راجع به چرایی پیشنهادها می‌شوند.

در این تحقیق برای تولید پیشنهادات، از یک مدل پیشنهاددهنده پایه مبتنی بر مدل‌سازی موضوع استفاده شد. اگر چه که مدل ارائه شده در این مقاله به تولید توضیحات تعقیبی، مستقل از الگوریتم پیشنهاددهنده می‌پردازد، اما در کارهای آتی به بررسی تاثیر انواع مدل‌های پیشنهاددهنده مختلف (مانند روش‌های پالایش همکارانه، روش‌های مبتنی بر محتوا و روش‌های ترکیبی) بر روی عملکرد مدل تولید توضیح پرداخته خواهد شد.

مراجع

- [1] S. de la Rouviere and K. Ehlers, "Lists as coping strategy for information overload on Twitter," presented at the Proceedings of the 22nd International Conference on World Wide Web, 2013, pp. 199–200.
- [2] V. Rakesh, D. Singh, B. Vinzamuri, and C. K. Reddy, "Personalized recommendation of twitter lists

- [15] G. Peake and J. Wang, "Explanation mining: Post hoc interpretability of latent factor models for recommendation systems," presented at the Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining, 2018, pp. 2060–2069.
- [16] C. Lonjarret, C. Robardet, M. Plantevit, R. Auburtin, and M. Atzmueller, "Why should i trust this item? explaining the recommendations of any model," presented at the 2020 IEEE 7th International Conference on Data Science and Advanced Analytics (DSAA), IEEE, 2020, pp. 526–535.
- [17] Y. Zhang, G. Lai, M. Zhang, Y. Zhang, Y. Liu, and S. Ma, "Explicit factor models for explainable recommendation based on phrase-level sentiment analysis," presented at the Proceedings of the 37th international ACM SIGIR conference on Research & development in information retrieval, 2014, pp. 83–92.
- [18] N. Wang, H. Wang, Y. Jia, and Y. Yin, "Explainable recommendation via multi-task learning in opinionated text data," presented at the The 41st International ACM SIGIR Conference on Research & Development in Information Retrieval, 2018, pp. 165–174.
- [19] H. Wang *et al.*, "Ripplenet: Propagating user preferences on the knowledge graph for recommender systems," presented at the Proceedings of the 27th ACM international conference on information and knowledge management, 2018, pp. 417–426.
- [20] Q. Ai, V. Azizi, X. Chen, and Y. Zhang, "Learning heterogeneous knowledge base embeddings for explainable recommendation," *Algorithms*, vol. 11, no. 9, p. 137, 2018.
- [21] Y. Xian, Z. Fu, S. Muthukrishnan, G. De Melo, and Y. Zhang, "Reinforcement knowledge graph reasoning for explainable recommendation," presented at the Proceedings of the 42nd international ACM SIGIR conference on research and development in information retrieval, 2019, pp. 285–294.
- [22] X. Wang, Q. Li, D. Yu, Q. Li, and G. Xu, "Reinforced path reasoning for counterfactual explainable recommendation," *IEEE Transactions on Knowledge and Data Engineering*, 2024.
- [23] F. Fusco, M. Vlachos, V. Vasileiadis, K. Wardatzky, and J. Schneider, "RecoNet: An Interpretable Neural Architecture for Recommender Systems.," presented at the IJCAI, 2019, pp. 2343–2349.
- [24] Y. Lu, R. Dong, and B. Smyth, "Why I like it: multi-task learning for recommendation and explanation," presented at the Proceedings of the 12th ACM Conference on Recommender Systems, 2018, pp. 4–12.
- [25] Z. Chen *et al.*, "Co-attentive multi-task learning for explainable recommendation.," presented at the IJCAI, 2019, pp. 2137–2143.
- [26] X. Wang, X. He, F. Feng, L. Nie, and T.-S. Chua, "Tem: Tree-enhanced embedding model for explainable recommendation," presented at the using content and network information," presented at the Proceedings of the International AAAI Conference on Web and Social Media, 2014, pp. 416–425.
- [3] L. Chen, Y. Zhao, S. Chen, H. Fang, C. Li, and M. Wang, "iplug: Personalized list recommendation in twitter," presented at the Web Information Systems Engineering–WISE 2013: 14th International Conference, Nanjing, China, October 13–15, 2013, Proceedings, Part II 14, Springer, 2013, pp. 88–103.
- [4] C.-H. Tsai and P. Brusilovsky, "The effects of controllability and explainability in a social recommender system," *User Modeling and User-Adapted Interaction*, vol. 31, pp. 591–627, 2021.
- [5] D. Shmaryahu, G. Shani, and B. Shapira, "Post-hoc Explanations for Complex Model Recommendations using Simple Methods.," presented at the IntRS@ RecSys, 2020, pp. 26–36.
- [6] Y. Zhang and X. Chen, "Explainable recommendation: A survey and new perspectives," *Foundations and Trends® in Information Retrieval*, vol. 14, no. 1, pp. 1–101, 2020.
- [7] V. Hassija *et al.*, "Interpreting black-box models: a review on explainable artificial intelligence," *Cognitive Computation*, vol. 16, no. 1, pp. 45–74, 2024.
- [8] I. Nunes and D. Jannach, "A systematic review and taxonomy of explanations in decision support and recommender systems," *User Modeling and User-Adapted Interaction*, vol. 27, pp. 393–444, 2017.
- [9] K. Balog and F. Radlinski, "Measuring recommendation explanation quality: The conflicting goals of explanations," presented at the Proceedings of the 43rd international ACM SIGIR conference on research and development in information retrieval, 2020, pp. 329–338.
- [10] R. Confalonieri, T. Weyde, T. R. Besold, and F. M. del Prado Martín, "Using ontologies to enhance human understandability of global post-hoc explanations of black-box models," *Artificial Intelligence*, vol. 296, p. 103471, 2021.
- [11] J. Schneider and J. Handali, "Personalized explanation in machine learning: A conceptualization," *arXiv preprint arXiv:1901.00770*, 2019.
- [12] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "Bert: Pre-training of deep bidirectional transformers for language understanding," *arXiv preprint arXiv:1810.04805*, 2018.
- [13] M. Caro-Martínez, G. Jiménez-Díaz, and J. A. Recio-García, "Conceptual modeling of explainable recommender systems: an ontological formalization to guide their design and development," *Journal of Artificial Intelligence Research*, vol. 71, pp. 557–589, 2021.
- [14] C. Nóbrega and L. Marinho, "Towards explaining recommendations through local surrogate models," presented at the Proceedings of the 34th ACM/SIGAPP Symposium on Applied Computing, 2019, pp. 1671–1678.

- Proceedings of the tenth ACM international conference on web search and data mining, 2017, pp. 485–494.
- [34] Y. Wu and M. Ester, “Flame: A probabilistic model combining aspect based opinion mining and collaborative filtering,” presented at the Proceedings of the eighth ACM international conference on web search and data mining, 2015, pp. 199–208.
- [35] Y. Zhang, “Incorporating phrase-level sentiment analysis on textual reviews for personalized recommendation,” presented at the Proceedings of the eighth ACM international conference on web search and data mining, 2015, pp. 435–440.
- [36] D. Kim, Y. Jo, I.-C. Moon, and A. Oh, “Analysis of twitter lists as a potential source for discovering latent characteristics of users,” presented at the ACM CHI workshop on microblogging, Citeseer, 2010.
- [37] C. Lu, W. Lam, and Y. Zhang, “Twitter user modeling and tweets recommendation based on wikipedia concept graph,” presented at the Workshops at the Twenty-Sixth AAAI conference on artificial intelligence, 2012.
- [38] R. Egger and J. Yu, “A topic modeling comparison between lda, nmf, top2vec, and bertopic to demystify twitter posts,” *Frontiers in sociology*, vol. 7, p. 886498, 2022.
- [39] M. Grootendorst, “BERTopic: Neural topic modeling with a class-based TF-IDF procedure,” *arXiv preprint arXiv:2203.05794*, 2022.
- Proceedings of the 2018 world wide web conference, 2018, pp. 1543–1552.
- [27] W. Ma *et al.*, “Jointly learning explainable rules for recommendation with knowledge graph,” presented at the The world wide web conference, 2019, pp. 1210–1221.
- [28] W. Sherchan, S. Nepal, and C. Paris, “A survey of trust in social networks,” *ACM Computing Surveys (CSUR)*, vol. 45, no. 4, pp. 1–33, 2013.
- [29] B. Wang, M. Ester, J. Bu, and D. Cai, “Who also likes it? generating the most persuasive social explanations in recommender systems,” presented at the Proceedings of the AAAI Conference on Artificial Intelligence, 2014.
- [30] C. Shi, Z. Zhang, P. Luo, P. S. Yu, Y. Yue, and B. Wu, “Semantic path based personalized recommendation on weighted heterogeneous information networks,” presented at the Proceedings of the 24th ACM International on Conference on Information and Knowledge Management, 2015, pp. 453–462.
- [31] Y. Zhang, X. Xu, H. Zhou, and Y. Zhang, “Distilling structured knowledge into embeddings for explainable and accurate recommendation,” presented at the Proceedings of the 13th international conference on web search and data mining, 2020, pp. 735–743.
- [32] J. Zheng, Z. Qin, S. Wang, and D. Li, “Attention-based explainable friend link prediction with heterogeneous context information,” *Information Sciences*, vol. 597, pp. 211–229, 2022.
- [33] Z. Ren, S. Liang, P. Li, S. Wang, and M. de Rijke, “Social collaborative viewpoint regression with explainable recommendations,” presented at the

Personalised Explanation Generation for Twitter List Recommendations using Semantic Similarity between Hashtags

Abstract: Twitter List recommender systems have achieved high prediction accuracy by leveraging diverse user and List information alongside complex algorithms. However, explainability remains a significant challenge in these systems. Providing meaningful explanations along with a set of recommendations can enhance user trust and satisfaction, assisting them in informed decision-making. In this paper, we present a model for the automated generation of personalized descriptions as explanations for recommended Twitter Lists. Specifically, our model selects frequently used hashtags from the content of the recommended List, establishing semantic relationships with the user's activity history. The aim is to present the List's subject in an understandable and personalized manner through a generated description. Through experiments conducted on a real Twitter dataset, our proposed model demonstrates its capability to generate explanations for a high percentage of the recommendations provided by a recommendation model.

Keywords: Explainable Recommender System, Personalized Explanation, Twitter List, Semantic Similarity between Hashtags