

تشخیص شایعات در شبکه های اجتماعی با استفاده از معماری ترکیبی LSTM - CNN و ارائه روش جدید پیش پردازش داده ها

مریم خسروی^۱، حسین شیرازی^{۲*}، کوروش داداش تبار^۳، سیدعلیرضا هاشمی گلپایگانی^۴

تاریخ دریافت: ۱۳۹۸/۱۲/۲۲

تاریخ پذیرش: ۱۳۹۹/۰۴/۲۴

چکیده

با توجه به جایگاهی که امروزه شبکه های اجتماعی در جوامع پیدا کرده اند، افراد بسیاری از این شبکه ها به منظور منتشر کردن عقاید و اطلاعات خود استفاده می کنند. یکی از چالش های امنیتی موجود در این شبکه ها، جلوگیری از حملات معنایی است. در حملات معنایی فرد مخرب قصد دارد تا با انتشار اطلاعات و شایعات نادرست در شبکه های اجتماعی، بر روی کاربران دیگر تاثیر بگذارد. بنابراین ایمنی افراد در این گونه شبکه ها به خطر می افتد. انتشار اطلاعات نادرست در مواقع بحرانی مانند جنگ یا انتخابات ممکن است، عواقب جبران ناپذیری برای یک اجتماع داشته باشد. از این رو در این پژوهش هدف اینست که بتوان شایعات از جمله شایعات فارسی را در شبکه های اجتماعی تشخیص داد. بدین منظور از یک معماری ترکیبی LSTM-CNN استفاده و برخلاف پژوهش های پیشین از نرخ یادگیری چرخشی بهره گرفته و روش جدیدی به منظور پردازش کردن داده ها قبل از ورود به شبکه برای بهبود نتایج ارائه شده است. علاوه بر آن نیز برای رفع مشکلات مربوط به کمبود داده مدل BERT برای تشخیص شایعات فارسی هم مورد بررسی قرار گرفت. در نهایت با ارزیابی روش پیشنهادی مشخص شد که این روش به دقت مناسبی برای تشخیص شایعات و همینطور شایعات فارسی تنها با بررسی محتوا، دست یافته است.

واژگان کلیدی: شبکه های اجتماعی، شایعه، یادگیری عمیق

^۱ دانشجوی دکتری، مجتمع دانشگاهی برق و کامپیوتر، دانشگاه صنعتی مالک اشتر، m.khosravi26@gmail.com

^{۲*} نویسنده مسئول، دانشیار، مجتمع دانشگاهی برق و کامپیوتر، دانشگاه صنعتی مالک اشتر، shirazi@mut.ac.ir

^۳ استادیار، مجتمع دانشگاهی برق و کامپیوتر، دانشگاه صنعتی مالک اشتر، Dadashtabar@yahoo.com

^۴ استادیار، دانشکده کامپیوتر و فناوری اطلاعات، دانشگاه صنعتی امیرکبیر، sa.hashemi@aut.ac.ir

۱. کلیات

با گسترش تعداد کاربران شبکه‌های اجتماعی برخط، این شبکه‌ها تبدیل به رسانه‌ی تاثیرگذاری به منظور به اشتراک گذاشتن اطلاعات شدند و افراد برای تصمیم‌گیری‌هایشان در مورد یک موضوع خاص مانند مباحث سیاسی، تحت تاثیر دوستان خود هستند. علاوه بر آن نیز شبکه‌های اجتماعی از جمله منابع مهم خبر برای بیشتر مردم محسوب می‌شوند. از این رو اطلاعات نادرست موجود در این شبکه‌ها، باعث تاثیر غلط در جمعیت زیادی از مردم و انتقال آلودگی‌های مختلف بین دوستان خواهد شد.

انتشار شایعه و اخبار نادرست از زمان‌های گذشته به عنوان یک مسئله‌ی اجتماعی و روانی با ابعاد گسترده بوده است که در مواقع بحرانی مانند جنگ بیشتر حائز اهمیت می‌شود. در هر زمان افراد و دولت‌ها ممکن است با اهداف مختلفی تصمیم به انتشار اخبار نادرستی بگیرند که از جمله این اهداف می‌توان به ایجاد بدبینی نسبت به نظام و مسئولان جامعه، افزایش نگرانی و اضطراب در بین مردم، ایجاد تقابل و صف بندی میان قشرهای مختلف و پاسخ‌گویی به حس انتقام، اشاره کرد.

در گذشته شایعات بیشتر از طریق گفت و گو، روزنامه‌ها، رادیو و تلویزیون منتشر می‌شد اما امروزه به دلیل قابلیت‌های موجود در شبکه‌های اجتماعی، این شبکه‌ها نیز به عنوان ابزاری برای انتشار شایعه تبدیل شده‌اند. شایعه به دلیل گیرایی و کششی که دارد با سرعتی باور نکردنی در جامعه حرکت می‌کند و گسترش آن می‌تواند تاثیرات مخرب زیادی را بر جوامع انسانی بگذارد و گاهی عواقب جبران‌ناپذیری را نیز به بار آورد.

با توجه به تاثیری که گسترش شایعات در جوامع از طریق شبکه‌های اجتماعی دارند، این موضوع و یافتن راهی به منظور مقابله با آن زمینه‌ی پژوهش‌های جدیدی در این حوزه شده است. با توجه به مسائل و مشکلات مطرح شده هنگام انتشار شایعات، در ایران نیز نیاز به وجود راهکاری به منظور تشخیص شایعه ضرورت دارد. از این رو در این پژوهش هدف اینست که روشی برای تشخیص شایعات در شبکه‌های اجتماعی با توجه به محتوا ارائه شود

بدین منظور ابتدا روش جدیدی به منظور پیش‌پردازش داده‌ها قبل از ورود به معماری یادگیری عمیق ارائه شده که باعث بهبود نتایج شده‌است و یادگیری ویژگی‌ها را ساده‌تر می‌کند.

علاوه بر آن برخلاف پژوهش‌های پیشین که نرخ یادگیری ثابت بود در این پژوهش از روش نرخ یادگیری چرخشی بهره می‌گیرد. در نهایت نیز مدل BERT را به منظور تشخیص شایعات فارسی مورد بررسی قرار می‌دهد و نتایج حاصل از آن ارزیابی خواهد شد.

در این پژوهش شایعه عبارت غیر معتبری فرض شده است که اطلاعاتی را بین مردم منتشر میکند. یک شایعه بعد از زمان خاصی می‌تواند با سه حالت درست (واقعی)، غلط (غیر واقعی) و نامشخص اتمام یابد. ویژگی مشترک همه‌ی آن‌ها ابهام، اهمیت و ایجاد اضطراب است. یک شایعه از یک یا چند منبع آغاز می‌شود و از گره‌ای به گره‌ی دیگر انتقال می‌یابد. ممکن است راجع به یک موضوع چندین شایعه‌ی مختلف درست یا نادرست وجود داشته باشد علاوه بر آن فرض اینست که محتوایی که کاربر منتشر می‌کند، متون دارای هشنگ است.

در ادامه ابتدا به بررسی کارهای انجام گرفته در این حوزه پرداخته خواهد شد و سپس راهکار پیشنهادی بیان می‌شود و در نهایت نتایج حاصل از این پژوهش مورد بررسی قرار خواهد گرفت.

۲. پیشینه

شایعات یکی از موضوعات علوم اجتماعی، در طول تاریخ بوده است تا این‌که در زمان جنگ جهانی دوم علاقه به روانشناسی شایعه و کنترل آن افزایش یافت [۱]. در آن زمان پژوهش اصلی توسط آلپورت و پستمن انجام گرفت و یکی از دلایل آن پژوهش این بود که روحیه‌ی مردم و امنیت جامعه توسط گسترش شایعات تهدید شده بود. در آن دهه مطالعات دیگری شکل گرفت و سپس این مطالعات خاموش شد. پس از آن نیز در اواخر دهه ۱۹۶۰ و ۱۹۷۰، به این موضوع گرایش ایجاد شد. در دو دهه‌ی اخیر نیز پژوهش‌های جدیدی در این زمینه انجام شده است [۲].

مطالعات انجام گرفته در این حوزه براساس ویژگی های مشترکی که دارند به چند دسته تقسیم بندی و بررسی می شوند. دسته اول مربوط به مطالعاتی هستند که انتشار شایعه را در اجتماع یا رسانه های اجتماعی از لحاظ روانشناختی مورد بررسی قرار داده اند. دسته دوم مربوط به مطالعات انجام شده ای امروزی است که از مدل های انتشار آلودگی برای انتشار شایعه بهره می گیرند و با استفاده از ویژگی های گراف انتشار شایعه و مدل ارائه شده، قادر به تشخیص منبع انتشار هستند که با مسدود کردن آن می توان از انتشار شایعه جلوگیری کرد. دسته آخر شایعات را با توجه به ویژگی های زبانشناختانه بررسی می کنند.

هر کدام از این دسته ها نیز به منظور رسیدن به هدفشان از روش خاصی بهره می گیرند که این روش ها را نیز می توان به سه دسته مهندسی ویژگی ها، یادگیری عمیق و ارائه ای مدل ریاضی انتشار تقسیم بندی کرد [۳].

همانگونه که گفته شد، دسته اول از لحاظ روانشناسی، شایعات را بررسی می کنند. مقاله ای [۴] از مفاهیم روانشناسی ادراکی به منظور بررسی گسترش اطلاعات نادرست به صورت عمد و غیر عمد در شبکه های اجتماعی استفاده می کند. هنگامی که فردی می خواهد درستی اطلاعات را ارزیابی کند، چهار معیار سازگار بودن با باورهای قبلی خود، منسجم بودن و متناقض نبودن پیام و بررسی منبع پیام را در نظر می گیرد. در مقاله ای [۵] نیز پژوهش های پیشین که در حوزه ای روانشناسی شایعه انجام گرفته بود و سعی به شناخت اختلافات کلیدی بین شایعات و غیر شایعات را داشت، بررسی شد. در واقع هدف اصلی این پژوهش، شناسایی الگوی انتشاری است که برای شایعات یکتاست [۶]. این دسته به تنهایی قادر به تشخیص نیست و برای رسیدن به هدف خود نیاز به استفاده از دو دسته - ی دیگر دارد.

دسته ی بعد، مدل های انتشار شایعه ای را ارائه می دهند و سپس روشی را برای واکنش کردن گره ها به منظور جلوگیری از انتشار شایعه معرفی می کنند. رساله ای [۷] معیار مرکزیت شبکه ای جدیدی به نام مرکزیت شایعه را معرفی کرده است و یک الگوریتم با زمان خطی به منظور محاسبه ای این معیار ارائه داده است در مقاله ای [۸] روشی برای یافتن منبع شایعه و ارزیابی احتمال اینکه قطعه ای از اطلاعات ممکن است شایعه باشد، مورد بررسی قرار گرفته است.

در رساله ای [۹] تئوری، الگوریتم ها و مدل های جدیدی برای فرایند انتشار در شبکه های پویا و ایستای بزرگ ارائه شده است. در رساله ای [۱۰] پویایی انتشار شایعه مورد مطالعه قرار گرفته است و هدف ارائه ای مدل جدیدی است که معیار پذیرش شایعه را در نظر بگیرد و هم چنین چگونگی واکنش کردن کارای گره ها در این شرایط به منظور توقف شایعات در شبکه های جهان کوچک و بی مقیاس مورد بررسی قرار گرفته است. در رساله ای [۱۱] هدف اینست که یک مدل انتشار اطلاعات جدید برای محیط های رسانه ای اجتماعی مبتنی بر نمایش همگن کاربر ارائه دهد و بتواند شایعات را از پست های مورد اعتماد تشخیص دهد. در مقاله ای [۱۲] هدف ارائه ای مدلی است که نقش حافظه، عقاید و تفاوت ها در رغبت به تولید تحریف و تفاوت ها در درجه ای اعتماد که مردم به یکدیگر دارند، بررسی شود. برخلاف مدل های قبلی، مدل ارائه شده قادر به مدلسازی واقعی تر از انتشار شایعه است. مشکل این دسته اینست که بعد از مدت زمانی که از انتشار شایعه گذشت، تشخیص صورت می گیرد که ممکن است این زمان طولانی باشد و شایعه اثر خود را گذاشته باشد.

در دسته ی سوم برای تشخیص شایعه از ویژگی های مرتبط با زبان و تحلیل محتوا استفاده شده است. در واقع در این پژوهش ها محتوای مطالب منتشر شده در رسانه های اجتماعی

مورد بررسی قرار می‌گیرد. این پژوهش نیز در این دسته قرار می‌گیرد. از این رو بیشتر راجع به این دسته بحث خواهد شد. در برخی پژوهش‌ها مانند مقالات [۱۳]، [۱۴] و [۱۵] در زبان فارسی از روش‌های سنتی تحلیل محتوا استفاده شده‌است. در این روش‌ها از پردازش زبان طبیعی به منظور وجود مجموعه‌ای از ویژگی‌ها استفاده و سپس از دسته‌بندی کننده‌های مختلف برای یافتن این مجموعه بهره گرفته می‌شد. با گسترش یادگیری عمیق دیگر نیازی به تعیین این ویژگی‌ها به صورت دستی وجود ندارد از این رو بیشتر پژوهش‌ها در دهه اخیر از آن به منظور بررسی محتوا استفاده می‌کنند.

یکی از اولین پژوهش‌هایی که از شبکه‌های عمیق به منظور تشخیص شایعه استفاده کرد، مقاله‌ی [۱۶] است. در این پژوهش از مدل مبتنی بر شبکه‌های عصبی بازگشتی به منظور تشخیص شایعات استفاده کرده است. [۱۷] نیز به منظور تشخیص شایعات دو مدل شبکه عصبی بازگشتی مبتنی بر ساختار درختی بالا به پایین و پایین به بالا ارائه داده‌است. در [۱۸] به بررسی شبکه‌های عصبی مختلف مانند شبکه عصبی کانولوشنال و شبکه عصبی بازگشتی به منظور تشخیص شایعات پرداخته است.

همان‌گونه که ذکر شد تاکنون پژوهش‌های زیادی در زمینه‌ی استفاده از یادگیری عمیق به منظور تشخیص شایعات انجام گرفته‌است اما تاکنون پیش پردازش هدفدار قبل از ورود داده به معماری یادگیری عمیق ارائه شده انجام نگرفته‌است.

در این پژوهش نیز به دلیل مزایای یادگیری عمیق از این روش برای بررسی محتوا بهره گرفته شده‌است اما برخلاف پژوهش‌های پیشین قبل از ورود داده به معماری مورد نظر روش پیش پردازش جدیدی ارائه شده‌است که بر روی داده‌ها اعمال و باعث بهبود نتایج می‌شود. علاوه بر آن کارکرد این روش‌ها تا کنون برای تشخیص شایعات فارسی مورد بررسی قرار نگرفته-

است. از این رو در این پژوهش مدل ارائه شده، بر روی مجموعه داده‌ی فارسی نیز ارزیابی شد.

۳. روش پیشنهادی

با گسترش روزافزون شبکه‌های اجتماعی برخط، حملات مختلف جدیدی هم در آن‌ها مطرح گردید. یکی از مهم‌ترین این حملات، امروزه انتشار اطلاعات نادرست و شایعه است.

انتشار اطلاعات نادرست ممکن است با اهداف مخرب گوناگونی مانند تخریب شخص یا دولت، ایجاد تنش و اضطراب در جامعه و گسترش بی اعتمادی صورت پذیرد که در مواقع بحرانی مثل جنگ و انتخابات می‌تواند آثار زیانبار و جبران ناپذیری را به اجتماع وارد سازد. به همین دلیل تشخیص شایعات در شبکه‌های اجتماعی، از اهمیت بالایی برخوردار است.

با توجه به اهمیت تشخیص شایعات، در این پژوهش از یادگیری عمیق بدین منظور استفاده شده‌است. زیرا یادگیری ویژگی‌ها در این روش به صورت خودکار انجام می‌گیرد و نیازی به تعریف آن‌ها به صورت دستی نیست.

علاوه بر آن در پژوهش‌های پیشین نرخ یادگیری استفاده شده برای معماری‌های ارائه شده به صورت ثابت در نظر گرفته می‌شد که مدت زمان آزمایش را افزایش و پیدا کردن نرخ مناسب را دشوار می‌ساخت که در این پژوهش به منظور حل این مشکل از نرخ یادگیری چرخشی استفاده شده‌است. به منظور بهبود نتایج نیز روش جدیدی برای پیش پردازش داده‌های ورودی به شبکه ارائه شده‌است که در ادامه به صورت کامل شرح داده خواهد شد.

۳-۱. روش مدل سازی کمی (ریاضی) مسئله

در این پژوهش شایعه عبارتی تعریف می‌شود که درستی آن تایید نشده یا غلط است. ویژگی مشترک همه‌ی آن‌ها ابهام، اهمیت و ایجاد اضطراب تعریف می‌شود. یک شایعه از یک یا چند منبع آغاز می‌شود و از گره‌ای به گره‌ی دیگر انتقال می-

تشخیص شایعات در شبکه های اجتماعی با استفاده از معماری....

متن فارسی مانند انواع مختلف نگارش برای یک کلمه، تنوع استفاده از "می" و "ها" چسبان و غیر چسبان، تنوع نگارش "ی" اضافه در کلمات مختوم به "ه" فواصل متفاوت میان کلمات را از بین برده و اصلاح می نماید. در واقع یکسان سازی هایی در این بخش اعمال می شود که در نتیجه آن، نظم مورد نظر در متن حکمفرما می شود [۱۹].

۲-۲-۳. روش پیش پردازش ارائه شده:

در این مرحله کلماتی که در شایعات و غیرشایعات مشترک هستند و پرکاربرد می باشند، شناسایی می شوند و سپس از هر دو مجموعه داده حذف می شود. این عملیات باعث می شود که کلماتی که برای تشخیص شایعات از غیر شایعات اهمیت کمتری دارند، حذف شوند و تمرکز معماری مورد نظر بر روی کلماتی که اهمیت بیشتری دارد، باشد و یادگیری بر این اساس انجام گیرد.

بدین منظور ابتدا کلمات پرکاربرد مجموعه داده ی شایعات را که با نام **R** و غیر شایعات که با نام **NR** نشان داده شده، مشخص می شوند. سپس بین این دو مجموعه اشتراک گرفته می شود. بنابراین کلماتی که در هر دو مجموعه مشترک هستند، جدا می شود. سپس از مجموعه داده ها این حذف صورت می گیرد. که به صورت معادلات ریاضی (۱) و (۲) در ادامه می آید:

$$S = \{NR \cap R\} \quad (1)$$

$$c_u = \{w'_1, w'_2, \dots, w'_n\} - \{n | n \in S\} \quad (2)$$

بنابراین این متن پردازش شده، به صورت $c_u = \{w'_1, w'_2, \dots, w'_n\}$ نشان داده می شود.

یابد. ممکن است راجع به یک موضوع چندین شایعه ی مختلف درست یا نادرست وجود داشته باشد.

در این پژوهش شبکه ی اجتماعی به صورت گراف $G=(V,E)$ در نظر گرفته می شود که **V** نشان دهنده ی گره ها (کاربران) هستند و **E** نشان دهنده ی روابط بین آن ها است. این کاربران توانایی انتشار متن دارند که این متن ممکن است توسط خود فرد تولید شده باشد یا اینکه آن را باز نشر داده باشد. هر فرد تعدادی دنبال کننده دارد که متونی را که منتشر می کند، می بینند و از طرف دیگر تعدادی افراد را دنبال می کند که قادر به دیدن متون انتشار داده شده توسط آن ها است. هر پست کاربر **u** با c_u نشان داده می شود و به صورت $c_u = \{w_1, w_2, \dots, w_n\}$ تعریف می شود که **n** تعداد کلمات هر پست را نشان می دهد. در نهایت نیز هدف اینست که مشخص شود c_u شایعه است یا خیر

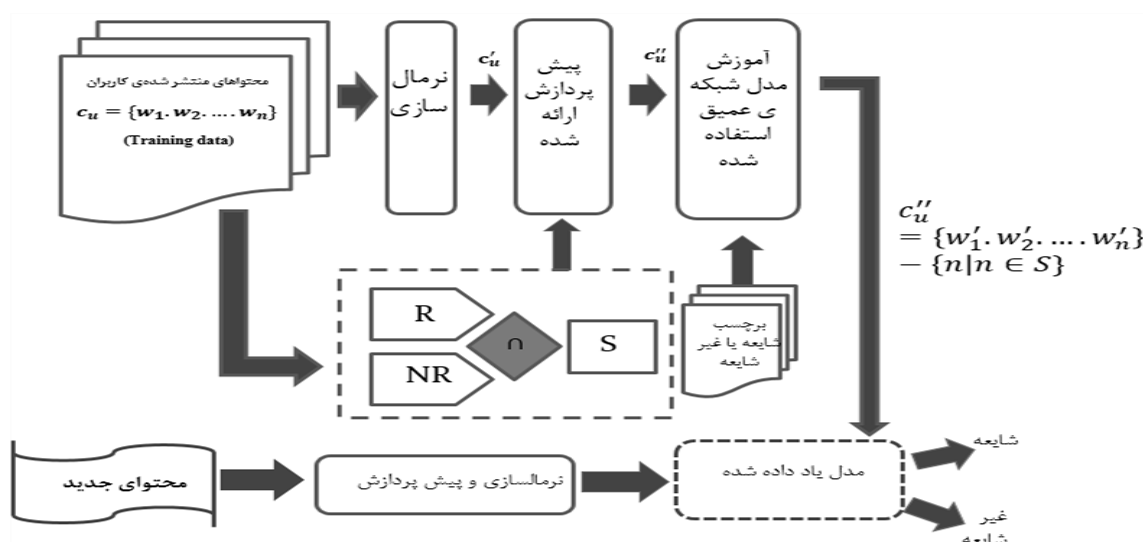
۲-۳. رویکرد و روش حل مسئله

همان گونه که گفته شد، هدف اینست که شایعه بودن یا نبودن متن مشخص شود. در واقع ورودی به صورت متن داده می شود و خروجی شایعه بودن یا نبودن متن را مشخص می کند، برای رسیدن به هدف ذکر شده نیاز است که ابتدا متن مورد نظر پیش پردازش و سپس از معماری ذکر شده استفاده شود. چگونگی عملکرد این روش در شکل ۱ نشان داده شده است. به منظور بهبود نتایج و جلوگیری از بیش برآزش و مشکل کمبود داده مدل **BERT**^۵ نیز برای داده ی فارسی مورد بررسی قرار گرفته است که در ادامه توضیح داده خواهد شد.

۳-۲-۱. پیش پردازش متن اولیه

از آن جایی که نتیجه ی این معماری بر روی متون فارسی بدون پیش پردازش مناسب نبود، برای دست یافتن به نتیجه ی دقیق تر باید ابتدا متن مورد نظر ویرایش شود بدین منظور از کتابخانه ی **Hazm** پایتون بهره گرفته شده است. پیش پردازش

⁵ Bidirectional Encoder Representation from Transformers



شکل ۱: مدل ارائه شده به منظور تشخیص شایعات

۳-۲-۳. استفاده از معماری LSTM-CNN

در این قسمت از مدل برداری کلمات فارسی تهیه شده در دانشگاه قم برای داده‌های فارسی استفاده شده است. سپس این تعبیه کلمه^۶ به لایه‌ی CNN و Pooling به منظور استخراج ویژگی‌های سطح بالاتر داده می‌شود.

سپس از یک لایه‌ی Dropout برای جلوگیری از بیش برازش استفاده می‌شود. از آنجایی که لایه‌ی LSTM برای داده‌هایی که به صورت دنباله‌ای هستند مانند جملات استفاده می‌شود، در اینجا هم خروجی به یک لایه‌ی LSTM وارد می‌شود و سپس به دو لایه‌ی متراکم با تابع relu برای کد کردن ویژگی‌ها داده می‌شود و پس از آن به یک لایه‌ی Dropout و در انتها به منظور دسته بندی نهایی به یک لایه‌ی متراکم با تابع sigmoid داده می‌شود. چگونگی معماری این قسمت در شکل ۲ آورده شده است. در ادامه هر قسمت به صورت کامل تر توضیح داده خواهد شد.

مدل برداری کلمات: هر کلمه با استفاده از روش Word2vec به صورت برداری نشان داده می‌شود [۲۰] که در این پژوهش از مدل برداری کلمات دانشگاه قم [۲۱] برای کلمات فارسی، استفاده شده است. بنابراین پس از اینکه کلمات به بردار تبدیل شدند هر متن به صورت $c_u = \{v_1, v_2, \dots, v_n\}$ نمایش داده می‌شوند.

لایه‌ی کانولوشن^۷: در این لایه، مجموعه‌ی m از فیلترها به یک پنجره‌ی لغزنده با طول h در طول جملات اعمال می‌شود. اگر $V[i:i+h]$ الحاق بردار کلمات v_i تا v_h را نشان دهد، ویژگی z_i برای یک فیلتر F به صورت معادله (۳) تولید می‌شود:

$$z_i = \sum_{k,j} (V_{i:i+h})_{k,j} \cdot F_{k,j} \quad (3)$$

الحاق همه‌ی بردارها در جمله، بردار ویژگی $z \in \mathbb{R}^{n-h+1}$ را تولید می‌کند. بردارهای z سپس در یک ماتریس نگاشت ویژگی $ZE \in \mathbb{R}^{m \times (n-h+1)}$ با هم تجمیع می‌شوند [۲۲].

Max Pooling: خروجی لایه‌ی کانولوشنال قبل از ورود به لایه‌ی Pooling از یک تابع فعالسازی غیرخطی می‌گذرد. در مرحله‌ی بعد المان‌های برداری با گرفتن ماکزیمم یک مجموعه ورودی ثابت غیر هم پوشان، تجمیع می‌شوند. ماتریس نگاشت ویژگی‌های تجمیع شده به صورت $Z_{pooled} \in \mathbb{R}^{m \times \frac{n-h+1}{s}}$ تعریف شده، به صورتی که s طول هر ورودی است. در صورتی که ورودی‌ها با مقدار گام s_t هم پوشان باشند، ماتریس نگاشت ویژگی به شکل $Z_{pooled} \in \mathbb{R}^{m \times \frac{n-h+1-s}{s_t}}$ تعریف می‌شود [۲۲].

⁷ convolutional neural network (CNN)

⁶ Word embedding

شده است و سپس دقت مدل بر روی مجموعه ی آزمایشی دیده نشده، مورد بررسی قرار گرفته است.

در این پژوهش بردار مربوط به ۱۰۰۰ کلمه ی پرتکرار مورد بررسی قرار می گیرد و بردارهای خروجی مربوط به هر کلمه نیز ۱۰۰ بعدی است، طول جملات نیز ۲۰۰ در نظر گرفته شده و سایدسته^{۱۳} ۳۲ است. از این رو ورودی لایه ی embedding به صورت (32,200) Embedding تعریف می-شود. خروجی آن نیز به صورت (۳۲,۲۰۰,۱۰۰) است که به عنوان ورودی لایه ی بعد یعنی لایه ی کانولوشن خواهد بود.

سپس خروجی به لایه ی بعد یعنی LSTM وارد می شود. البته در طول بررسی مشخص شد که استفاده از تعبیه کلمات از پیش یادگیری شده ی دانشگاه قم کارایی LSTM را بالا می برد. LSTM نیز با ۱۲۸ نورون نتیجه ی مناسب تری ایجاد کرد. سپس دو لایه ی Dense با تعداد ۶۴ نورون تعریف شده است و در انتها نیز از sigmoid برای مشخص کردن کلاس شایعه استفاده می شود. تعداد epoch نیز ۱۰۰ برای داده های فارسی تعیین شد.

• استفاده از نرخ یادگیری چرخشی

یکی از مهم ترین ابر پارامترها تعیین مقدار نرخ یادگیری بهینه است. روشی که در مقاله ی [۲۴] به منظور حل این مشکل استفاده شده است، نرخ یادگیری چرخشی نام دارد که نیاز به آزمایش های فراوان برای تعیین نرخ بهینه را از بین می برد. در واقع در این روش به صورت چرخشی این نرخ را بین مقادیر محدوده تغییر می دهد. کمترین و بیشترین مقدار نرخ یادگیری هم از طریق آزمایش محدوده ی نرخ یادگیری^{۱۴} مشخص می شود. که پس از آزمایش محدوده بین (۰,۰۰۱)، (۰,۰۰۰۱) تعیین شد. این روش به دلیل تغییر نرخ یادگیری از بیش برآزش^{۱۵} جلوگیری می کند و دقت را بهبود می بخشد. علاوه بر آن به دلیل اینکه نرخ را تغییر می دهد نیاز به بررسی همه ی نرخ های یادگیری نیست و با تعداد تکرار ۲۵ که حدود ۱/۴ قبل است می توان به نتیجه ی بهتری رسید. به همین دلیل

لایه ی LSTM: لایه ی LSTM^۸ بهبود یافته ی شبکه ی عصبی بازگشتی^۹ است که در واقع با معرفی بلوک های حافظه، مشکلات RNN را رفع کرده است. دنباله ی ورودی $z = (z_1, z_2, \dots)$ به معماری LSTM داده می شود. در این معماری سه گیت ورودی، خروجی و فراموشی به منظور محاسبه ی خروجی h_t ، در سلول حافظه به صورت تکراری از $t=1$ تا T در لایه ی پنهان بازگشتی LSTM مطابق معادلات (۴) تا (۹) تعریف می شود [۲۳].

$$f_t = \sigma(W_f \times z_t + U_f \times h_{t-1} + b_f) \quad (4)$$

$$i_t = \sigma(W_i \times z_t + U_i \times h_{t-1} + b_i) \quad (5)$$

$$C_t = \tanh(W_c \times z_t + U_c \times h_{t-1} + b_c) \quad (6)$$

$$C_t = i_t \times C_t + f_t \times C_{t-1} \quad (7)$$

$$o_t = \sigma(W_o \times z_t + U_o \times h_{t-1} + b_o) \quad (8)$$

$$h_t = o_t \times \tanh(C_t) \quad (9)$$

در این معادلات W ماتریس های وزندهی و b بایاسها می باشند.

لایه ی پنهان: سپس از دو لایه ی متراکم^{۱۰} مطابق معادله ی (۱۰) با تابع فعال سازی relu استفاده شده است و خروجی به یک لایه ی متراکم با تابع فعال سازی sigmoid مطابق معادله ی (۱۱) برای دستیابی به نتیجه ی نهایی L_i داده شده است که در واقع نشان می دهد ورودی متن C_i شایعه بوده است یا خیر.

$$P_i = \text{Relu}(W^T h + b) \quad (10)$$

$$L_i = \text{sigmoid}(W^T P + b) \quad (11)$$

۳-۲-۴. تنظیم کردن ابر پارامترها

به منظور میزان کردن ابر پارامترها^{۱۱} نیز از روش اعتبارسنجی متقابل ۳-تایی^{۱۲} و اعتبار سنجی hold-out استفاده

⁸ Long short term memory

⁹ Recurrent Neural Network (RNN)

¹⁰ Dense

¹¹ Hyper-parameter

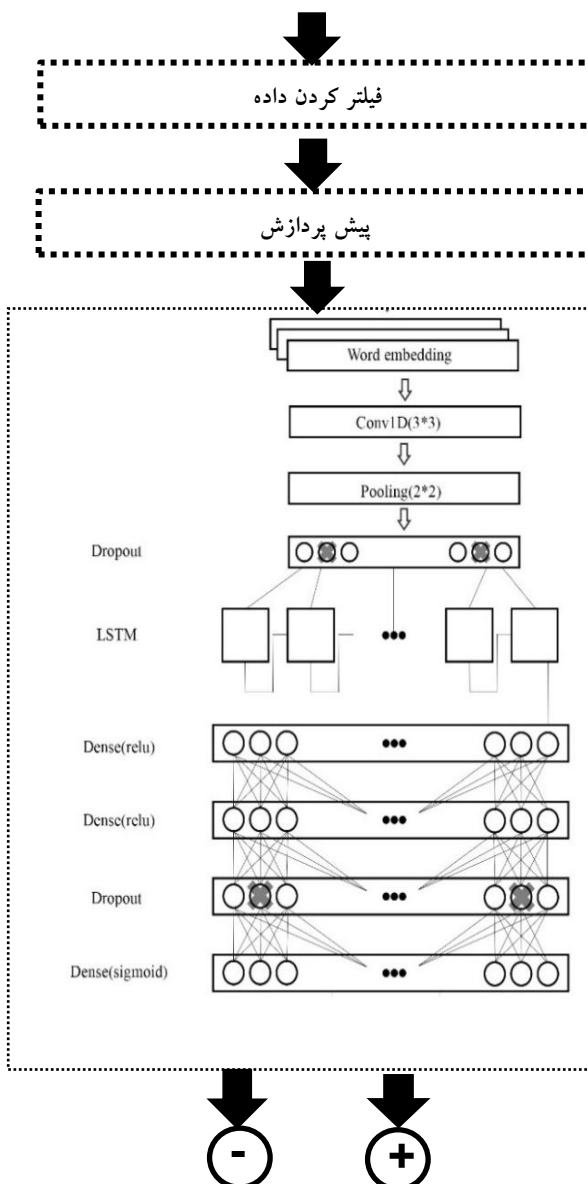
¹² 3-fold cross validation

¹³ Batch

¹⁴ LR range test

¹⁵ Overfit

در معماری LSTM-CNN از این روش به منظور بررسی بهبود نتایج بهره گرفته شده است.



شکل ۲. معماری استفاده شده برای تشخیص شایعه بودن متن

مقداردهی اولیه شده و سپس همه‌ی پارامترها با استفاده از داده‌های برچسب دار، وظایف فروسو^{۱۸} تعدیل می‌شوند.

هر وظیفه‌ی فروسو، مدل تعدیل شده‌ی جداگانه‌ای دارد، حتی اگر با مقادیر یکسان مقداردهی اولیه شده باشند. در این مدل برخلاف مدل‌های قدیمی از مدل‌های زبانی راست به چپ یا چپ به راست استفاده نشده است بلکه از وظایف نظارت نشده به منظور پیش آموزش BERT استفاده شده است. بدین منظور درصدی از ورودی به صورت تصادفی پوشانده می‌شود و سپس نیاز است که آن‌ها حدس زده شوند. سپس بردار پنهان مربوط به توکن‌های حدس زده شده به عنوان خروجی لایه‌ی softmax کلمات داده می‌شوند.

از آنجایی که بسیاری از وظایف فروسو مانند پرسش و پاسخ و استنتاج زبان طبیعی مبتنی بر فهم ارتباط بین جملات هستند، به منظور یاد دادن مدلی که ارتباط بین جملات را بفهمد، در این مقاله یک پیش آموزش برای پیش بینی جمله‌ی بعدی به صورت باینری انجام شده است. به این صورت که هنگام انتخاب جمله‌ی A و B برای هر مثال، ۵۰ درصد امکان دارد که B جمله‌ی بعدی باشد.

در این پژوهش از مدل Bert-Base, multilingual cased استفاده شده است که یک مدل ۱۲ لایه با ۱۱۰ میلیون پارامتر شامل ۱۰۴ زبان مختلف است که فارسی را نیز شامل می‌شود و بر روی داده‌های ویکی پدیا یادگیری انجام شده است [۲۶]. ورودی داده‌های مربوط به این پژوهش است و خروجی نیز لایه‌ای اضافه شده که بتواند مشخص کند داده‌ی مورد نظر شایعه است یا خیر.

۴. ارزیابی

در این فصل به بررسی و ارزیابی روش ارائه شده به منظور بررسی محتوا پرداخته خواهد شد. برای بررسی روش‌های یادگیری عمیق به منظور رسیدن به هدفی که در این مقاله دنبال می‌شود، چهار معیار صحت، دقت، فراخوانی و F-score ارزیابی می‌شوند. از این رو در این پژوهش نیز همین معیارها مورد بررسی قرار می‌گیرد.

در [۲۵] یک مدل بازنمایی زبان جدید با نام BERT ارائه شده که شامل دو مرحله‌ی پیش آموزش^{۱۶} و تعدیل کردن^{۱۷} است. در مرحله‌ی پیش آموزش، مدل بر روی داده‌های بدون برچسب آموزش داده می‌شود. برای تعدیل کردن نیز مدل BERT ابتدا با پارامترهای پیش آموزش داده شده‌ی وظایف،

^{۱۶} Pre-train

^{۱۷} Fine-tune

^{۱۸} Downstream tasks

۴-۱. مجموعه داده

در این پژوهش بررسی بر روی مجموعه داده ای استفاده شده در مقاله ی [۱۵] انجام شده که شامل حدود ۲۰۰۰ پست توئیت است که تعدادی از آن ها شایعه می باشند. این مجموعه داده شامل ۷۸۳ شایعه جمع آوری شده از سایت های گمانه و شایعات است. به همین تعداد توئیت غیر شایعه نیز به داده اضافه شده است. به منظور حفظ ویژگی تصادفی بودن، پست ها از زمانی جمع آوری شده اند که اولین پست شایعه منتشر شده است.

از این مجموعه داده تنها متن پست ها در این پژوهش مورد بررسی قرار گرفته اند که حدود ۶۰ درصد از پست های شایعه و ۳۰ درصد پست های غیر شایعه با زبان محاوره هستند. مشخصات این مجموعه داده در جدول ۱ آمده است.

به دلیل تعداد کم این داده به منظور بهبود ارزیابی از روش افزایش داده^{۱۹} به نام **EDA**^{۲۰} در مقاله ی [۲۷] استفاده شده است. در این روش از ۴ عملیات ساده اما قدرتمند به منظور افزایش داده بهره گرفته شده که عبارتند از: جایگذاری با لغات هم معنی، حذف، درج و جابه جایی به صورت تصادفی.

جدول ۱. مشخصات داده ای استفاده شده برای پژوهش

شایعات	غیر شایعات	
درصد ۳۰	درصد ۷۰	درصد محاوره بودن تعداد متون
۱۴۰	۱۴۰	حداکثر تعداد کلمات
۴	۴	حداقل تعداد کلمات
۷۸۳	۱۳۰۰	تعداد
۳۷۱۷	۲۶۷۳	تعداد داده ها بعد از افزایش داده

به منظور ارزیابی این روش برای محتوای انگلیسی نیز از مجموعه داده **PHEME** [۲۸] استفاده شده است. این مجموعه داده از شایعات و غیر شایعات مربوط به پنج اتفاق مختلف جمع آوری شده است که برای هر اتفاق پوشه ای شامل این دو دسته وجود دارد. در این پژوهش از همه ی مجموعه داده استفاده شده است.

۴-۱-۴. داده های آموزشی و آزمایشی

داده های مورد استفاده در این پژوهش ابتدا به دو دسته ی آموزشی و آزمایشی تقسیم می شوند که به نسبت ۸۰ به ۲۰ می باشند. در واقع روش مورد نظر بر روی داده های آزمایشی دیده نشده بررسی خواهد شد.

به منظور جلوگیری از بیش برازش و اطمینان از نتایج آموزش، ۲۰ درصد از داده های آموزشی نیز به صورت تصادفی به عنوان داده های اعتبارسنجی انتخاب می شود.

۴-۲. یافته ها

در این قسمت در شش حالت بررسی شایعه بودن متن انجام و سپس ارزیابی صورت گرفت:

حالت اول: استفاده از معماری **LSTM-CNN** و **LSTM** با تغییر نرخ یادگیری و بررسی نتایج برای مجموعه داده فارسی

حالت دوم: استفاده از نرخ یادگیری چرخشی برای معماری **LSTM-CNN** و **LSTM** برای مجموعه داده فارسی

حالت سوم: استفاده از روش افزوده سازی داده ها و بررسی تغییرات برای مجموعه داده فارسی

حالت چهارم: بررسی روش ارائه شده برای پیش پردازش و فیلتر کردن داده ها برای مجموعه داده فارسی

حالت پنجم: استفاده از مدل **BERT** برای مجموعه داده فارسی

حالت ششم: بررسی معماری **LSTM-CNN** و روش پیش پردازش ارائه شده برای مجموعه داده انگلیسی

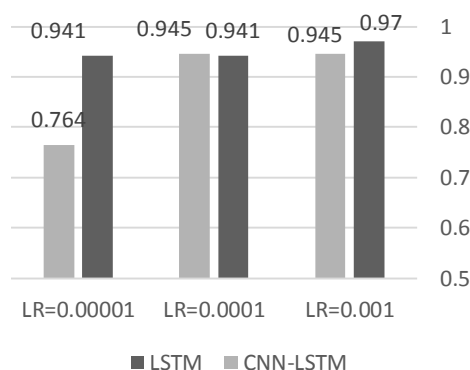
سپس با بررسی نتایج هر یک مشخص شد که در حالت دوم بدون نیاز به آزمایش همه ی حالات به منظور یافتن نرخ یادگیری بهینه و بدون اعمال هزینه ی اضافه به معماری، می تواند نرخ یادگیری بهینه را پیدا و نتیجه ی مناسبی را ایجاد کند. در حالت سوم که داده ها افزایش می یافت دقت بالاتر می رفت و از بیش برازش جلوگیری می شد.

در حالت چهارم با استفاده از این پیش پردازش نتایج بهبود یافت و نسبت به سه حالت قبل بهتر شد.

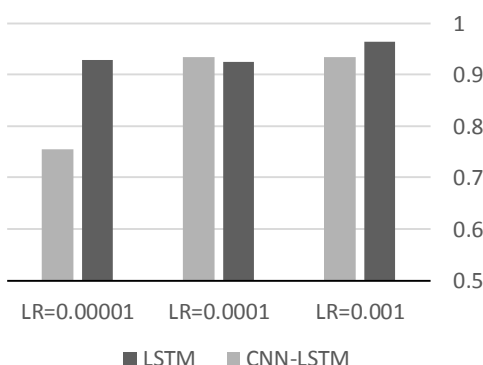
در حالت پنجم نیز که از مدل **BERT** استفاده شد به دلیل اینکه این مدل ابتدا بر روی داده های ویکی پدیا آموزش دیده، مشکل کمبود داده را تا حد زیادی جبران و نتایج بهتری نسبت به حالت اول تولید می کند اما از لحاظ دقت مانند حالت دوم

¹⁹ Data Augmentation

²⁰ Easy Data Augmentation



شکل ۳: مقایسه دقت دو معماری LSTM و LSTM-CNN با نرخ‌های یادگیری مختلف بر روی داده‌های آزمایشی فارسی



شکل ۴: مقایسه F-Score دو معماری LSTM و LSTM+CNN با نرخ‌های یادگیری مختلف بر روی داده‌های آزمایشی فارسی با توجه به نمودارها مشخص می‌شود که بهترین حالت، استفاده از نرخ یادگیری با مقدار ۰,۰۰۱ و معماری LSTM است.

حالت دوم بررسی، استفاده از روش نرخ یادگیری چرخشی بود که نتایج در جدول ۲ آمده است. با توجه به نتایج جدول و نمودارهای قبل، بهترین نتیجه زمانی رخ می‌دهد که از معماری CNN-LSTM با نرخ یادگیری چرخشی استفاده شود.

جدول ۲: ارزیابی حالت دوم در صورت استفاده از نرخ یادگیری چرخشی برای داده‌های فارسی آزمایشی

فراخوانی	صحت	f-score	دقت	معیار ارزیابی
۰,۹۷۵	۰,۹۵۶	۰,۹۶۵	۰,۹۷۰	LSTM-clr
۰,۹۸۰	۰,۹۶۲	۰,۹۷۴	۰,۹۷۸	CNN-LSTM-clr

است. بهترین نتیجه مربوط به حالت چهارم است از این پیش پردازش برای مدل BERT نیز می‌توان استفاده کرد.

سپس به بررسی نتایج حاصل از این پژوهش و مقاله‌ی [۱۵] پرداخته می‌شود. روشی که در مقاله‌ی [۱۵] پیشنهاد شده است اگرچه مبتنی بر اطلاعات شبکه و محتوا است اما بیشتر براساس اطلاعات شبکه پیش بینی شایعات را انجام می‌دهد و اگر به تنهایی مورد استفاده قرار گیرد، قادر به تشخیص شایعه نخواهد بود.

معماری پیشنهادی در این رساله تنها محتوا را در نظر می‌گیرد و قادر است با دقت بالایی شایعات را تشخیص دهد. در واقع دقتی بالاتر از مقاله‌ی ذکر شده را ایجاد می‌کند بدون آن که به بررسی سایر ویژگی‌های شبکه بپردازد.

علاوه بر آن برای مجموعه داده‌ی انگلیسی نیز با مقاله‌ی [۱۸] مقایسه انجام گرفته است که در قسمت بعد به صورت کامل توضیح داده خواهد شد.

۲-۴-۱. نتایج ارزیابی یافته‌ها

همان‌گونه که گفته شده شش حالت در این پژوهش مورد بررسی قرار گرفت که نتایج حاصل از بررسی آن‌ها و همینطور مقایسه‌ی نتایج حاصل از این روش با مقالات قبل در این قسمت آورده می‌شود. لازم به ذکر است که این بررسی بر روی داده‌های آزمایشی دیده نشده پس از آموزش مدل، صورت گرفته است.

حالت اول بررسی دقت و F-Score بر روی مجموعه داده - ی فارسی آزمایشی با استفاده از معماری LSTM+CNN و LSTM با نرخ‌های یادگیری مختلف است. مقایسه دقت و F-Score به ترتیب در شکل ۳ و شکل ۴ آورده شده‌اند.

تشخیص شایعات در شبکه های اجتماعی با استفاده از معماری....

حالت سوم استفاده از افزوده سازی داده ها است که همان-
طور که در جدول (۳) آمده نتایج را به مقدار کمی بهبود
بخشیده است.

جدول (۳): ارزیابی حالت سوم بعد از افزوده سازی داده ها برای داده-

های فارسی آزمایشی

فراخوانی	صحت	f-score	دقت	معیار ارزیابی
۰,۹۸۰	۰,۹۶۹	۰,۹۷۴	۰,۹۷۸	CNN-LSTM-aug

حالت چهارم نیز به بررسی روش پیش پردازش پیشنهادی
می پردازد. همان طور که در جدول ۴ مشخص است نتایج به
صورت خیلی واضحی بهبود یافته است.

جدول ۴: ارزیابی حالت چهارم بعد از اعمال روش پیش پردازش برای

داده های فارسی آزمایشی

فراخوانی	صحت	f-score	دقت	معیار ارزیابی
۰,۹۹۴	۰,۹۹۴	۰,۹۹۳	۰,۹۹۱	CNN-LSTM

حالت پنجم نیز استفاده از مدل BERT بود که همان گونه
که نتایج در جدول (۵) مشخص است. دقتی نزدیک به روش
پیشنهادی یعنی CNN-LSTM با نرخ چرخشی، دارد.

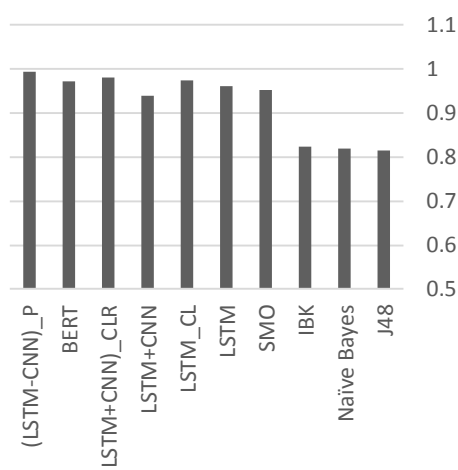
جدول (۵): نتایج دقت حاصل از استفاده از BERT برای داده های

فارسی آزمایشی

صحت	فراخوانی	f-score	دقت	معیار
۰,۹۸۳	۰,۹۷۲	۰,۹۷۸	۰,۹۷۷۶	BERT

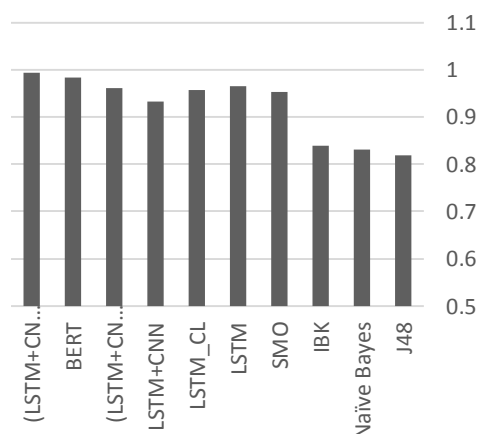
نتایج صحت و فراخوانی استفاده از روش ذکر شده در این
پژوهش، در مقایسه با مقاله ی [۱۵] به ترتیب در شکل ۵ و
شکل ۶ آورده شده است.

در [۱۵] تشخیص شایعه با استفاده از روش های یادگیری
ماشین J48، SMO، IBK و NaiveBayes انجام گرفته است.
البته لازم به ذکر است که نتایج دقت در این روش ها با در نظر
گرفتن ویژگی های شبکه نیز هست در حالی که در روش ذکر
شده در این پژوهش تنها بررسی بر روی محتوا انجام گرفته
شده است.



شکل ۵: مقایسه ی صحت روش های این پژوهش با روش های مقاله-

ی [۱۵]



شکل ۶: مقایسه ی فراخوانی روش های این پژوهش با روش های

مقاله ی [۱۵]

با توجه به نمودار مشخص می شود که روش های یادگیری
عمیق پیاده سازی در این پژوهش با وجود اینکه تنها محتوا را
مورد بررسی قرار می دهد، از همه ی حالت های بررسی ویژگی-
های انتشار در روش های یادگیری ماشین نظارت شده بیان
شده در مقاله ی [۱۵]، کارایی بالاتری دارد و حتی در حالتی که
در این مقاله از هر دو ویژگی متنی و انتشار استفاده کرده است
نیز بهتر عمل می کند و تنها در روش SMO است که کارایی
پایین تری نسبت به برخی از آن ها دارد که آن هم چون در
روش های یادگیری عمیق تنها محتوا مورد بررسی قرار گرفته
است بازهم نتیجه ی قابل قبولی است.

علاوه بر آن، یکی از روش‌های ذکر شده در مقاله‌ی [۱۵] به عنوان نمونه به نام روش **IBK** پیاده سازی شد و به منظور بررسی ویژگی‌های ذکر شده‌ی مربوط به محتوا در مقاله، تنها با این ویژگی‌ها روش گفته شده با معماری مورد نظر مقایسه شد و نتیجه‌ی صحت و دقت در جدول ۶ مورد مقایسه قرار گرفته‌است.

جدول ۶: مقایسه‌ی بررسی دسته بندی کننده‌ی **IBK** با توجه به ویژگی‌های محتوایی [۱۵] با روش **(LSTM_CNN)_CL** و **(LSTM_CNN)_P**

صحت	دقت	
۰,۹۶۲	۰,۹۷۸	(LSTM_CNN)_CL
۰,۵۰۲	۰,۵۰۵	مقاله‌ی [۱۵]
۰,۹۹۴	۰,۹۹۱	(LSTM_CNN)_P

حالت ششم نیز به بررسی روش ارائه شده و کارایی آن بر روی مجموعه داده انگلیسی مرتبط است. در این قسمت نتایج با مقاله‌ی [۱۸] مقایسه شده‌است. که در مقاله مقدار صحت و فراخوانی محاسبه نشده بود. از این رو معماری **CNN-LSTM** مجدد مورد ارزیابی قرار گرفته و نتایج در جدول ۷ آمده‌است. با توجه به جدول مشخص می‌شود که با استفاده از این روش به دلیل تمرکز بر روی کلمات دارای اهمیت بیشتر، نتایج بهبود می‌یابند.

جدول ۷: مقایسه‌ی روش ارائه شده برای مجموعه داده انگلیسی و نتایج مقاله‌ی [۱۸]

فراخوانی	صحت	f-score	دقت	معیار ارزیابی
-	-	۴۹,۲۳	۶۶,۷۶	CNN-LSTM[18]
۷۴,۰۶	۷۴,۲	۷۴,۱۲	۷۴,۴۷	CNN-LSTM پیاده سازی شده
۷۶,۳۳	۷۶,۰۴	۷۵,۹۵	۷۶,۰۴	(CNN_LSTM)_P

۵. جمع بندی

امروزه با گسترش شبکه‌های اجتماعی و استفاده‌ی اقشار مختلف مردم از این شبکه‌ها و امکان انتشار محتوا در آن‌ها، بستر جدیدی برای سوء استفاده‌ی برخی افراد نیز فراهم

گردیده‌است. برخی افراد با استفاده از امکانات شبکه‌های اجتماعی با اهداف گوناگون سعی در انتشار اطلاعاتی می‌کنند که یا درستی آن‌ها هنوز اثبات نشده‌است و یا اینکه نادرست است و به این صورت باعث ایجاد تنش و نابسامانی در جامعه را ایجاد می‌کنند. از این رو تشخیص این گونه اطلاعات و جلوگیری از انتشار آن‌ها اهمیت زیادی پیدا کرده است.

در دهه‌های اخیر پژوهش‌های زیادی در این حوزه با استفاده از یادگیری عمیق انجام گرفته‌است اما هیچ کدام پیش پردازش هدفداری به منظور بهبود نتایج ارائه نداده‌اند. علاوه بر آن تاکنون این روش‌ها برای تشخیص شایعات فارسی در شبکه‌های اجتماعی، مورد بررسی قرار نگرفته‌اند.

در روش‌های یادگیری عمیق پیشنهاد شده برای تشخیص شایعه، در سایر زبان‌ها نیز نرخ یادگیری ثابت در نظر گرفته می‌شد که هزینه‌ی اجرا را زیاد می‌کرد در این پژوهش با توجه به نتایج حاصل از استفاده از روش نرخ یادگیری چرخشی در معماری **LSTM_CNN** برای تشخیص شایعه در متون فارسی، مشخص می‌شود این روش، نتایج بهتری را تولید می‌کند.

در این پژوهش موارد زیر مورد بررسی قرار گرفت:

- ارائه‌ی روش جدیدی به منظور پیش پردازش داده
- یافتن معماری یادگیری عمیق مناسب مبتنی بر **LSTM-CNN** برای تشخیص شایعات فارسی
- استفاده از نرخ یادگیری چرخشی به جای نرخ یادگیری ثابت به منظور صرف هزینه‌ی کمتر و دستیابی به نتیجه‌ی بهتر.
- بررسی مدل **BERT** برای شایعات فارسی و ارزیابی نتایج آن

یکی از مشکلاتی که در این روش وجود دارد نیاز به حجم زیاد داده است که روش‌های مختلفی از جمله افزودن داده‌ها برای حل آن وجود دارد که یک نمونه از آن در این پژوهش استفاده شد. روش‌های متفاوت دیگر در آینده می‌تواند مورد بررسی قرار گیرد.

علاوه بر آن معماری‌های یادگیری عمیق دیگری وجود دارند که ممکن است نتایج بهتری را تولید کنند که نیازمند بررسی می‌باشند. روش پیش پردازش ارائه شده نیز می‌تواند بهبود یابد که این مورد را نیز باید مورد ارزیابی قرار داد.

۶. منابع

- [12]Wang,Chao,Xuan Tan,Zong, Ye,Ye,Wang,Lu, Cheong,K.H and Xie,Neng-gang,(2017), A rumor spreading model based on information entropy, Scientific Reports,pp.1-14
- [13]Hamidian,S and Diab,M,(2019) Rumor Detection and Classification for Twitter Data, SOTICS, The Fifth International Conference on Social Media Technologies, Communication, and Informatics,pp.1-7
- [14]Zhang,Q,Zhang,S,Dong,J,Xiong,J, and Cheng,X,(2015)Automatic Dectong of Rumor on Social Network, in Natural Language Processing and Chinese Computing, pp.113-122
- [15]Zamani,Somayeh,Asadpour,Masoud,Moazzami,Dar a,(2017)Rumor Detection for Persian Tweets,the Iranian Conference on Electrical Engineering (ICEE), pp.1532-1536
- [16]Ma,Jing,Gao,Wei,Mitra,Prasenjit,Kwon,Sejeong,Bernard,J.J,(2016),Detecting rumors from microblogs with recurrent neural networks,Proceedings of the Twenty-Fifth International Joint Conference on Artificial Intelligence,pp.3818-3824.
- [17]Asghar,M.Z,Habib,Ammara,Habib,Anam,Khan,A, Ali,R and Khattak.A,(2019) Exploring deep neural networks for rumor detection, Journal of Ambient Intelligence and Humanized Computing,pp.1-19
- [18]Ma,J, Gao.W, and Wong,K.F,(2018),Rumor Detection on Twitter with Tree-structured Recursive Neural Networks, [Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics](#),pp.1980-1989
- [19] Pavithra C P,S.J, (2019), Deep Learning Approach For Rumour Detection In Twitter: A Comparative Analysis, the International Conference on Systems, Energy & Environment (ICSEE),pp.1-5
- [۲۰] سپهریان، زهرا، سدیدپور، سعیده سادات، شیرازی، حسین، (۲۰۱۴)، روشی جدید در خلاصه سازی متون فارسی بر اساس عبارت پرس و جوی کاربر، پدافند الکترونیکی و سایبری، صص ۵۲-۶۳
- [21]Mikolov,T, Sutskever,I, Chen,K, Corrado,G and Dean,J,(2013), Distributed representations of words and phrases and their compositionality, the 26th International Conference on Neural Information Processing Systems - Volume 2, pp.1-9.
- [۲۲] شاقوزی، محمدرضا(۹۵/۱۱/۲۰)، مدل برداری کلمات فارسی در ویکی پدیا، قابل دسترسی در: <http://dml.qom.ac.ir/tag/word2vec>
- [1]Tanaka,Yuko, Sakamoto, Yasuaki and Matsuka,T,(2012), Transmission of Rumor and Criticism in Twitter after the Great Japan Earthquake, Annual Meeting of the Cognitive Science Society, pp. 2387-2392.
- [2]Sunstein,Cass.R,(2009) On Rumors How Falsehoods Spread, Why We Believe Them, What Can Be Done. Princeton University Press publisher, Farrar, Straus and Giroux.pp. 1-90
- [3]Ben Veyseh,Amir Pouran,Thi My.T,Thien, Huu Nguyen and Dou,Dejing,(2019),Rumor Detection in Social Networks via Deep Contextual Modeling, IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining,pp.113-120
- [4]Kumar,K P Krishna and Geethakumari,G,(2014), Detecting misinformation in online social networks using cognitive psychology, Human-centric Computing and Information Sciences, vol. 4, no. 1, pp. 1-22
- [5]Kwon, Sejeong , Cha, Meeyoung , Jung,K, Chen,W and Wang,Y, (2013),Aspects of Rumor Spreading on a Microblog Network, in Social Informatics, vol. 8238, (Lecture Notes in Computer Science: Springer International Publishing, pp. 299-308.
- [6]Allport,G.W, Postman,Leo,(1947)The psychology of rumor, Journal of Clinical Psychology, New York: Russell & Russel,pp.1-247 .
- [7]Zaman,T.R, (2013)"Information Extraction with Network Centralities: Finding Rumor Sources, Measuring Influence, and Learning Community Structure", PhD Thesis. Massachusetts Institute of Technology,Department of Electrical Engineering and Computer.
- [8]Seo,E,Mohapatra,P and Abdelzaher,T,(2012) Identifying rumors and their sources in social networks, SPIE, pp. 1-13
- [9]Prakash,B.A,(2012)Understanding and Managing Propagation on Large Networks—Theory, Algorithms, and Models,PhD Thesis,School of Computer Science Carnegie Mellon University,Computer Science Department
- [10]SINGH.A,(2013),Rumor Dynamics and Inoculation of Nodes in Complex Networks, PhD Thesis, Indian Institute Of Technology Kanpur, Department Of Electrical Engineering
- [11]Liu,Yang,Xu,Songhua,(2016),Detecting Rumors Through Modeling Information Propagation Networks in a Social Media Environment, IEEE Transactions on Computational Social Systems, vol. 3, pp. 46-62

[23]Deriu,Jan and Cieliebak,M(2016), Sentiment Analysis using Convolutional Neural Networks with Multi-Task Training and Distant Supervision on Italian Tweet, Proceedings of Third Italian Conference on Computational Linguistics (CLiC-it 2016) & Fifth Evaluation Campaign of Natural Language Processing and Speech Tools for Italian. Final Workshop (EVALITA 2016),pp.1-5.

[24]Rehman,A.U,Malik,A.K,Raza,B and Ali,W (2019), A Hybrid CNN-LSTM Model for Improving Accuracy of Movie Reviews Sentiment Analysis, Multimedia Tools and Applications, vol. 78, no. 18, pp. 26597-26613.

[25]Smith,L.N(2017), Cyclical Learning Rates for Training Neural Networks, 2017 IEEE Winter Conference on Applications of Computer Vision (WACV), pp. 464-472.

[26]Devlin,Jacob, Lee,Kenton,Toutanova,Kristina (2018)BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding, Proceedings of NAACL-HLT 2019, pp.4171-4176

[27] Devlin.J. (2019). google bert multilingual. Available:
<https://github.com/googleresearch/bert/blob/master/multilingual.md>

[28]Wei, Jason, Zou, Kai (2019), EDA:Easy Data Augmentation Techniques for Boosting Performance on Text Classification Tasks, EMNLP/IJCNLP,pp.6382-6386.

[29]Kochkina,Elena,Liakata,Maria,Zubiaga,Arkatiz(10.06.2018) , PHEME dataset for Rumour Detection and Veracity Classification, Availabel:
https://figshare.com/articles/PHEME_dataset_for_Rumour_Detection_and_Veracity_Classification/6392078