

Providing a Three-Dimensional Tensor Approach For Classifying and Detecting Fake News - A Case Study of Persian News in The Field of COVID-19¹

Vahid Mottaghi

PhD. Candidate in IT Management, Department of IT Management, Qeshm Branch, Islamic Azad University, Qeshm, Iran. Mvahid500@gmail.com

Mahdi Esmaili

Assistant Professor, Department of Computer Science, Kashan Branch, Islamic Azad University, Kashan, Iran (**Corresponding Author**). m.esmaeili@iaukashan.ac.ir

Ghasem Ali Bazaei

Assistant Professor, Department of Management, Central Tehran Branch, Islamic Azad University, Tehran, Iran. Gh.bazaei@iauctb.ac.ir

Mohammad Ali Afshar Kazemi

Associate Professor, Department of Management, Central Tehran Branch, Islamic Azad University, Tehran, Iran. M_afsharkazemi@iauec.ac.ir

Abstract

Purpose: Convolutional neural networks, as one of the most important models of deep learning, have gained high accuracy on these issues. In this study, discussion and text analysis at the sentence level and improving the performance of neural networks to detect fake news has been convolution. The network of words for bags of words in the data model so that each word according to the two-dimensional vector space to become matrices. One of the limitations of convolutional networks is that it works at the word level and cannot consider the relationship and distance between sentences. And sentence-level analysis is a major problem in this research. Sentence level analysis is a major problem in this research.

1. **This present study is taken from:** PhD. Thesis, Information Technology Management, business Intelligence Orientation, Student: Vahid Mottaghi, **Improving deep learning approaches to the problem of detecting fake news A Case Study of Persian News in The Field of COVID-19.** Supervisor: Mahdi Esmaili, Ghasem Ali Bazaei, Advisor: Mohammad Ali Afshar Kazemi, Islamic Azad University, Qeshm Branch.

Received: 2021/06/05 ; **Accepted:** 2021/08/07

Mottaghi, Vahid; Esmaili, Mahdi; Bazaei, Ghasem Ali & Afshar Kazemi, Mohammad Ali (2021). Providing a Three-Dimensional Tensor Approach For Classifying and Detecting Fake News - A Case Study of Persian News in The Field of COVID-19. *Sciences and Techniques of Information Management*, 7(4). DOI: 10.22091/stim.2021.7014.1592

© the authors / **Publisher:** University of Qom



In this research, a basic model based on convolutional networks is proposed in which documents are given to the network in the form of 3D tensors to solve the mentioned problem. Considering 3D tensors allows the model to learn the position of words in a sentence and achieve more accurate results in detecting fake news.

Methodology: This study is applied research in which about 42,000 Persian news from different cities of Iran were collected from Twitter and using preprocessing, additional and useless data is deleted and after tagging the deleted texts, the news text is used for the proposed approach using Python software and related libraries.

Findings: During testing, some machine learning algorithms had more power in classification problems, but with the changes in the structure of the convolutional network algorithm, better results were obtained than machine learning algorithms and other similar algorithms.

Conclusion: Considering 3D tensors allows the model to learn the position of words in a sentence, and this proposed model has gained considerable accuracy compared to the proposed approaches in the literature. The proposed model without adding additional overhead in terms of the number of features and network depth, by changing the input has been able to achieve better and more acceptable results than other approaches in the literature and achieve an accuracy of more than 94%.

Keywords: Natural Language Processing, Text Classification, Convolutional Neural Networks, Fake News Detection.

ارائه رویکرد تنسور سه بعدی برای طبقه‌بندی و تشخیص اخبار جعلی: مطالعه موردی اخبار فارسی در حوزه کرونا و ویروس^۱

وحید متقی

دانشجوی دکتری، گروه مدیریت فناوری اطلاعات، واحد قشم، دانشگاه آزاد اسلامی، قشم، ایران.

Mvahid500@gmail.com

مهدی اسماعیلی

استادیار، گروه علوم کامپیوتر، واحد کاشان، دانشگاه آزاد اسلامی، کاشان، ایران (نویسنده مسئول).

m.esmaeili@iaukashan.ac.ir

قاسمعلی بازایی

استادیار، گروه مدیریت، واحد تهران مرکزی، دانشگاه آزاد اسلامی، تهران، ایران. Gh.bazaei@iauctb.ac.ir

محمدعلی افشار کاظمی

دانشیار، گروه مدیریت، واحد تهران مرکزی، دانشگاه آزاد اسلامی، تهران، ایران. M_afsharkazemi@iauec.ac.ir

چکیده

هدف: هدف پژوهش حاضر اختصاص یکی از کلاس‌های جعل و واقعی به متن‌های آزاد می‌باشد. شبکه‌های عصبی کانولوشنی به عنوان یکی از مهم‌ترین مدل‌های یادگیری عمیق، دقت بالایی را بر روی این مسائل بدست آورده است. در این تحقیق آنالیز متن در سطح جمله و بهبود عملکرد شبکه عصبی کانولوشنی جهت تشخیص اخبار جعلی مورد توجه بوده است. در این شبکه‌ها کلمات به صورت کیسه‌ای از کلمات به مدل داده می‌شوند که هر کلمه با توجه به فضای برداری به ماتریس‌های دو بعدی تبدیل می‌شود. یکی از محدودیت‌های شبکه‌های کانولوشن این است که در سطح کلمه کار کرده و نمی‌تواند رابطه و فاصله بین جملات را در نظر بگیرد و آنالیز در سطح جمله مشکل اساسی در این تحقیق می‌باشد. در این پژوهش یک مدل پایه‌ای مبتنی بر شبکه‌های کانولوشنی پیشنهاد شده که در آن اسناد به صورت تنسورهای سه بعدی به شبکه داده می‌شوند تا بتواند مشکل مذکور را مرتفع نماید. در نظر گرفتن تنسورهای سه بعدی امکان یادگیری موقعیت کلمات در جمله را برای مدل فراهم می‌آورد و به نتایج دقیق‌تری در تشخیص اخبار جعل دست می‌یابد.

روش‌شناسی: پژوهش حاضر مطالعه‌ای کاربردی بوده که در آن حدود ۴۲۰۰۰ اخبار فارسی از شهرهای مختلف ایران از توییتر جمع‌آوری شده و با عمل پیش‌پردازش، داده‌های اضافی و غیر مفید حذف و پس از برچسب زدن متون پاک‌سازی شده، متن اخبار جهت رویکرد پیشنهادی با استفاده از نرم‌افزار پایتون پردازش شده‌اند.

۱. **پژوهش حاضر برگرفته از:** رساله دکتری، رشته مدیریت فناوری اطلاعات، گرایش کسب و کار هوشمند، دانشجو: وحید متقی، با عنوان: **بهبود رویکردهای یادگیری عمیق برای مسأله تشخیص اخبار جعلی: مطالعه موردی اخبار فارسی حوزه کرونا و ویروس**، استاد راهنما: مهدی اسماعیلی، قاسمعلی بازایی، استاد مشاور: محمدعلی افشار کاظمی، ارائه شده در دانشگاه آزاد اسلامی - واحد قشم است.

تاریخ دریافت: ۱۴۰۰/۰۳/۱۵ ؛ تاریخ پذیرش: ۱۴۰۰/۰۵/۱۶

یافته‌ها: برخی از الگوریتم‌های یادگیری ماشین دارای قدرت بیشتری در مسائل طبقه‌بندی بودند، ولی با تغییراتی که در ساختار الگوریتم شبکه کانولوشن صورت گرفت، نتایج بهتری نسبت به الگوریتم‌های یادگیری ماشین و سایر الگوریتم‌های مشابه حاصل شد.

نتیجه‌گیری: در نظر گرفتن تنسورهای سه بعدی امکان یادگیری موقعیت کلمات در جمله را برای مدل فراهم می‌آورد و این مدل پیشنهادی در مقایسه با رویکردهای پیشنهادی در ادبیات، دقت قابل توجهی را بدست آورده است. مدل پیشنهادی بدون اضافه کردن سربار اضافی از لحاظ تعداد ویژگی‌ها و عمق شبکه، با تغییر در ورودی توانسته است به نتایج بهتر و قابل قبول از سایر رویکردهای موجود در ادبیات دست یافته و به دقت و صحت بیش از ۹۴ درصد دست یابد.

کلیدواژه‌ها: پردازش زبان طبیعی، طبقه‌بندی متن، شبکه‌های عصبی کانولوشنی، تنسور سه بعدی، اخبار جعلی، اخبار فارسی، کرونا ویروس.

۱. مقدمه

امروزه نمی‌توانیم زندگی خود را بدون دسترسی به شبکه جهانی وب تصور کنیم. هر کس از اینترنت برای مقاصد مختلفی، برای مثال جهت جستجو و یا ارسال پیام یا انتشار خبر استفاده می‌کند. اطلاعات به راحتی می‌توانند توسط کاربران در وبلاگ‌ها، انجمن‌ها و شبکه‌های اجتماعی منتشر شوند. ظهور اخبار جعلی در انتخابات ریاست جمهوری ایالات متحده در سال ۲۰۱۶ نه تنها خطرات ناشی از اخبار جعلی، بلکه چالش‌های مطرح شده هنگام تلاش برای جدا کردن اخبار جعلی از اخبار واقعی را نشان داد (درچنسکی^۱ و بونچوا^۲، ۲۰۱۴). اخبار جعلی ممکن است یک اصطلاح نسبتاً جدید باشد، اما لزوماً پدیده جدیدی نیست (اندروفر^۳، ۲۰۱۷). با این حال، امروزه پیشرفت‌های فناوری و انتشار اخبار از طریق رسانه‌های مختلف باعث گسترش اخبار جعلی شده است. به همین ترتیب، نرخ اخبار جعلی اخیراً به صورت تصاعدی افزایش یافته و لازم است اقداماتی صورت گیرد تا از گسترش آن در آینده جلوگیری شود. کشف اخبار جعلی و طبقه‌بندی اخبار با توجه به صحت تشخیص آن براساس اخبار جعل یا واقعی است. در یک نگاه ساده، این یک کار طبقه‌بندی دو دویی می‌باشد، در حالی که در یک محیط متفاوت، این یک کار طبقه‌بندی چندگانه نیز به شمار می‌آید (ژاو^۴ و زعفرانی^۵، ۲۰۱۸).

وقایع سیاسی اخیر منجر به محبوبیت و گسترش اخبار جعلی شده است. برخی از این اخبارها به گونه‌ای گسترش پیدا می‌کنند که به واقعیت بسیار نزدیک هستند و بسیاری از افراد را تحت تاثیر قرار می‌دهند. روش‌های مختلفی برای خودکارسازی روند کشف اخبار جعلی ارائه شده‌اند که از محبوب‌ترین روش‌ها می‌توان به لیست‌های سیاه از منابع و نویسندگان غیرقابل اعتماد اشاره کرد. در حالی که این ابزارها مفید هستند، برای ایجاد یک راه‌حل کامل‌تر برای کشف چنین اخباری، ما باید موارد بیشتری را در نظر بگیریم. روش‌های یادگیری ماشین و تکنیک‌های پردازش زبان طبیعی، ابزارهای مفیدی برای تشخیص اخبار جعلی هستند (ژاو و زعفرانی، ۲۰۱۸؛ شو^۶ و همکاران،

<http://stjm.gom.ac.ir>

1. Derczynski
2. Bontcheva
3. Andorfer
4. Zhou
5. Zafarani
6. Shu

۲۰۱۷). روش‌های یادگیری عمیق، به دلیل انتخاب ویژگی‌ها به صورت خودکار در تشخیص اخبار جعلی به شدت مورد استفاده قرار می‌گیرند.

در مورد شبکه‌های عصبی کانولوشنی^۱ که از الگوریتم‌های یادگیری عمیق بوده، انتخاب ویژگی‌ها به صورت خودکار می‌باشد. این شبکه‌ها در سطح کلمه کار می‌کنند و نمی‌توانند رابطه و فاصله بین جملات را در نظر بگیرند. برای این منظور نیاز است که این رابطه‌ها در نظر گرفته شود، چرا که ممکن است این رابطه‌ها دارای ویژگی‌های مفیدی باشند. ایده کلیدی این رویکرد استفاده از اطلاعات مکانی هر جمله در اسناد می‌باشد. شایان ذکر است که در نظر گرفتن جملات در کنار هم می‌تواند اطلاعات اضافه‌ای را ایجاد کند که در مدل‌های مبتنی بر کیسه‌ای از کلمات^۲ این اطلاعات قابل دسترس نیستند.

در تحقیق پیش‌رو، برای حل مشکل آنالیز در سطح کلمه، مجموعه‌ای از تانسورهای سه بُعدی پیشنهاد شده است. هر یک از اسناد ورودی، به یک تانسور سه بُعدی تبدیل می‌شود. در بعد اول سند، بعد دوم کلمات و بعد سوم بردارهای تعبیه شده قرار می‌گیرند که در آن هر سند ورودی به دو حد آستانه لایه‌گذاری^۳ می‌شود که برای این منظور از لایه‌گذاری صفر^۴ استفاده می‌شود. در این رویکرد لازم است که تغییرات گسترده‌ای در لایه‌های شبکه کانولوشنی انجام شود تا مدل بتواند تانسورهای سه بُعدی را کنترل کند.

۲. چارچوب نظری

مشکل اخبار جعلی موضوع جدیدی نیست و می‌توان استدلال کرد که از ابتدا با گسترش روزنامه‌های کاغذی وجود داشته است. انسان‌ها همواره تمایل دارند که اطلاعات غلط را باور کنند. تشخیص دادن اخبار جعلی برای انسان‌ها دشوار است. می‌توان ادعا کرد که تنها راه شناسایی آنها در صورت عدم داشتن داده‌های زیاد، به صورت دستی است (ژاو و زعفرانی، ۲۰۱۸).

اکنون اخبار جعلی به عنوان یکی از بزرگ‌ترین تهدیدات برای دموکراسی، روزنامه‌نگاری و آزادی بیان تلقی می‌شود و اعتماد عمومی را به دولت‌ها تضعیف می‌کند. تأثیر احتمالی آن در

1. Convolutional neural network

2. Bag Of Words

۳. تانسور به مجموعه‌ای از داده‌ها در یک دسته گفته می‌شود برای مثال بردار یک تانسور یک بُعدی و ماتریس یک تانسور دو بُعدی است.

4. Padding

5. Zero padding

همه‌پرسی بحث‌برانگیز "Brexit" و انتخابات ریاست جمهوری آمریکا در سال ۲۰۱۶ مورد توجه قرار گرفت (شارما^۱ و همکاران، ۲۰۱۹). اقتصاد ما نیز از انتشار خبرهای جعلی مصون نمی‌باشد، زیرا اخبار جعلی به نوسانات بورس سهام و معاملات گسترده مرتبط است. به عنوان مثال، اخبار جعلی که ادعا می‌کرد باراک اوباما در اثر انفجار زخمی شده است، ۱۳۰ میلیارد دلار ارزش سهام را از بین برد (پوگ^۲، ۲۰۱۷). این وقایع و ضررها باعث ایجاد انگیزه در تحقیقات خبری جعلی شده و بحث پیرامون اخبار جعلی را برانگیخته است.

به عنوان یک بستر ایده‌آل برای تسریع در انتشار اخبار جعلی، رسانه‌های اجتماعی فاصله فیزیکی بین افراد را می‌شکنند و بسترهای غنی را برای اشتراک نظرات و بررسی فراهم کرده و کاربران را به مشارکت و بحث در مورد اخبار آنلاین تشویق می‌کنند (ژاو و زعفرانی، ۲۰۱۸).

هیچ تعریف جهانی واحدی برای اخبار جعلی حتی در علم روزنامه‌نگاری وجود ندارد. تعریف واضح و دقیق به ایجاد یک پایه محکم برای تجزیه و تحلیل اخبار جعلی و ارزیابی مطالعات مرتبط کمک می‌کند. مطالعات موجود غالباً اخبار جعلی را با اصطلاحات و مفاهیمی مانند اخبار کذب^۳ (سیرینگ^۴، کاک^۵ و دیوکار^۶، ۲۰۱۶؛ راپازا^۷، ۲۰۱۷)، اخبار دروغین^۸، اخبار طنز^۹ (الکات^{۱۰} و جنتسکاو^{۱۱}، ۲۰۱۷)، اطلاعات غلط^{۱۲} (شو و همکاران، ۲۰۱۷)، اطلاعات نادرست^{۱۳} (والدراپ^{۱۴}، ۲۰۱۷) و شایعات^{۱۵} (وئوقی^{۱۶}، روی^{۱۷} و آرال^{۱۸}، ۲۰۱۸) مرتبط می‌سازند.

1. Sharma
2. Pogue
3. Maliciously false news
4. Siering
5. Koch
6. Deokar
7. Rapoza
8. False news
9. Satire news
10. Allcott
11. Gentzkow
12. Disinformation
13. Misinformation
14. Waldrop
15. Rumor
16. Vosoughi
17. Roy
18. Aral

بر این اساس اخبار جعلی و تعاریف آن در جدول ۱ به طور خلاصه آمده است.

جدول ۱- تفاوت بین تعاریف مختلف اخبار جعل

نوع	اعتبار	قصد	اخبار
اخبار کذب	خیر	مخرب	بله
اخبار دروغین	خیر	غیر معلوم	بله
اخبار طنز	غیر معلوم	غیر مخرب	بله
اطلاعات غلط	خیر	مخرب	غیر معلوم
اطلاعات نادرست	خیر	غیر معلوم	غیر معلوم
شایعات	غیر معلوم	غیر معلوم	غیر معلوم

بدیهی است که رایج ترین اصطلاحات در مراجع اصلی، اخبار جعلی و شایعات هستند، اما با این وجود محققان جنبه های دیگری را که مربوط به اطلاعات غلط در وب هستند، از قبیل clickbait، اسپم های اجتماعی و بررسی های جعلی را نیز مورد تجزیه و تحلیل قرار داده اند. در ادبیات، طبقه بندی های مختلفی از اخبار جعلی و شایعات ارائه شده است که بیشتر به منبع و نوع داده های مورد استفاده برای تجزیه و تحلیل بستگی دارد. مطالعات اولیه در این زمینه، به ویژه از منظر محاسباتی، در چند سال اخیر انجام شده است. بنابراین، مرزهای مطالعه اخبار جعلی اغلب به روشنی مشخص نیست. به همین دلیل، ما معتقدیم که ارائه بینش در مورد اینکه چه نوع داده ای می تواند تبدیل به موضوع تجزیه و تحلیل و چگونگی تعریف آن شود، اساسی است. در ادامه دو نوع از این اطلاعات آمده است:

اخبار جعلی: واژه اخبار جعلی برای شناسایی اطلاعات نادرست در رسانه های اصلی، به ویژه برای مطالب مرتبط با وب، مطرح شده است. بنابراین، اخبار جعلی مقالاتی خبری هستند که عمداً برای گمراه کردن یا سوءاستفاده از خوانندگان نوشته شده اند، اما می توان آنها را با استفاده از منابع دیگر تایید کرد. چندین مطالعه اخیر، مانند برکویتز^۱ و شوارتز^۲ (۲۰۱۶)، تعریف اخبار جعلی بیان شده را تایید کرده اند.

در پژوهش کشتی^۳ و وواس^۴ (۲۰۱۷) تمایز میان جنبه های مختلف اخبار جعلی معرفی شده

1. Berkowitz
2. Schwartz
3. Kshetri
4. Voas

است. به طور خاص، نویسندگان بر جعل جدی، جعل در مقیاس بزرگ و تقلب‌های طنز توجه کرده‌اند. جعل‌های جدی نمونه اولیه اخبار جعلی هستند، یعنی مقالاتی با قصد مخرب می‌باشند (مثلاً مصاحبه‌های جعلی، مقالات شبه علمی و غیره) که اغلب از طریق رسانه‌های اجتماعی به صورت ویروس منتشر می‌شوند. جعل در مقیاس بزرگ، گزارشی از اطلاعات دروغین است که به عنوان خبرهای مناسب پنهان شده‌اند (کشتی و و آس، ۲۰۱۷). معمولاً چنین فریبکاری‌هایی در مقیاس بزرگ‌تر از یک مقاله خبری ساده سازماندهی می‌شوند و اغلب افراد یا ایده‌های عمومی را هدف قرار می‌دهند. جعل‌های طنزآمیز به منظور سرگرم کردن خوانندگان، که به نظر، آنها از هدف طنز نویسنده آگاه است، نوشته شده است. نمونه‌هایی از قطعات طنز به عنوان اخبار واقعی، مانند مواردی که توسط وب سایت lercio.it تولید می‌شوند، در دسترس هستند.

سه وجه اصلی اخبار جعلی قابل شناسایی است: شکل آن، به عنوان مقاله خبری، قصد فریبنده آن، که می‌تواند هم طنزآمیز باشد و هم بدخواه، اثبات مطالب آن، که به صورت کامل یا جزئی نادرست است.

شایعات: در ادبیات علمی اخیر، شایعات احتمالاً بیشترین اطلاعات دروغ مورد مطالعه در وب بوده‌اند. آنها به اطلاعاتی اشاره می‌کنند که هنوز توسط منابع رسمی تایید نشده‌اند و بیشتر توسط کاربران در سیستم عامل‌های رسانه‌های اجتماعی پخش شده‌اند.

شایعات محصول عصر اینترنت نیستند، با مطالعات اولیه مربوط به پایان جنگ جهانی دوم می‌توان دریافت که همواره شایعات وجود داشته‌اند (کچارسکی^۱، ۲۰۱۶؛ کانروی^۲، رابین^۳ و چن^۴، ۲۰۱۵). با این حال، می‌توان استدلال کرد که اینترنت و بسترهای رسانه‌های اجتماعی به طور خاص زمینه باروری برای انتشار اطلاعات غیرقابل اثبات و شایعات هستند (رابین، چن و کانروی، ۲۰۱۵).

ارائه یک تعریف رسمی از شایعه ساده نیست. در حقیقت، محققان تفسیرهای مختلفی را گزارش کرده‌اند. یکی از تعاریف جذاب و پذیرفته شده توسط نویسندگان در پژوهش آلپورت^۵ و

1. Kucharski
2. Conroy
3. Rubin
4. Chen
5. Allport

پستمن^۱ (۱۹۴۶) آمده است. آنها در تحقیقات خود، شایعات را به عنوان بیانیه‌های اطلاعاتی تایید نشده و در گردش معرفی می‌کنند. علاوه بر این، میر^۲ (۱۹۶۹) شایعه را به طور خاص به عنوان داستان در حال گردش که معتبر و مشکوک است، تعریف می‌کند، که ظاهراً معتبر است، اما تایید آن دشوار بوده و باعث شک و تردید و یا اضطراب می‌شود. برای این تعریف، یک شایعه باید یک واکنش تأثیرگذار بر مخاطبان خود ایجاد کند. با این حال، می‌توان استدلال کرد که این تعاریف به ویژگی تایید نشدن اطلاعات وابسته هستند. این اطلاعات تایید نشده می‌تواند درست، تا حدی درست، کاملاً نادرست بوده یا تایید نشده باقی بمانند (میر، ۱۹۶۹).

سرانجام، مطالعات مختلف سعی کرده‌اند شایعات را با توجه به نوع، دامنه و ویژگی‌های آنها طبقه‌بندی کنند. به عنوان مثال، ناپ^۳ (۱۹۴۴) شایعات را با توجه به واکنش مورد انتظار از منظر روانشناسی طبقه‌بندی کرده است. از دیدگاه عملی‌تر، زوبیگا^۴ و همکاران (۲۰۱۸) شایعات جالب را به دو دسته اصلی، شایعات طولانی مدت نشان‌دهنده اطلاعاتی که برای مدت طولانی گردش می‌کند و ممکن است هیچ‌گاه صحت یا نادرستی آنها تایید نشوند و شایعات خبری متداول که می‌توانند محصول اطلاعات غلط غیر عمدی باشند، تقسیم کردند. شایعات خبری توجه بیشتری به ادبیات با توجه به شایعات طولانی داشته است. این در شرایطی است که چنین شایعاتی می‌توانند در یک دوره زمانی کوتاه خطرناک‌تر باشند و برای جلوگیری از گسترش آنها، باید هرچه سریع‌تر شناسایی شوند، به ویژه اگر قصد آنها مخرب باشد.

۳. پیشینه پژوهش

در چند سال گذشته، علاقه جامعه پژوهش به اخبار و شایعات جعلی به میزان قابل توجهی افزایش یافته است. شکل ۱ تعداد مقالات مربوط به شایعات و اخبار جعلی منتشر شده از سال ۲۰۰۶ تا ۲۰۱۸ را نشان داده و در پایگاه داده Scopus^۵ طبقه‌بندی شده است. می‌توان مشاهده کرد که علاقه به شایعات در طی دوازده سال گذشته به طور پیوسته در حال رشد بوده است، در حالی که اخبار جعلی عمدتاً در سه سال گذشته مورد توجه عظیمی قرار گرفته‌اند. شکل ۲ تارنمایی از توزیع

1. Postman

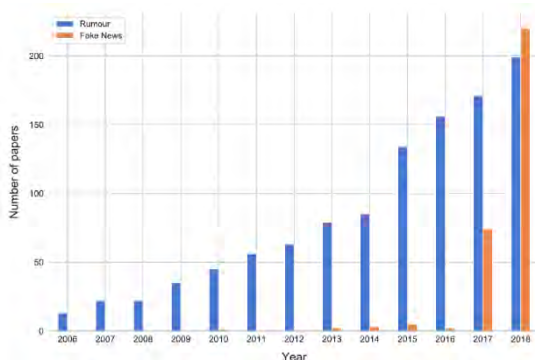
2. Meyer

3. Knapp

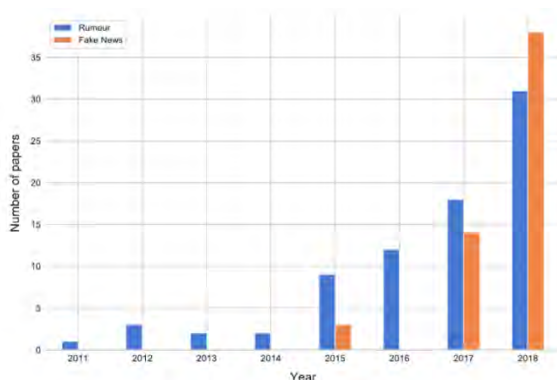
4. Zubiaga

5. <https://www.scopus.com/>

مقالات چاپ شده در ۱۱ سال اخیر می‌باشد.



شکل ۱- توزیع چاپ مقالات در حوزه اخبار جعلی و تشخیص شایعات (زوبیگا و همکاران، ۲۰۱۵)

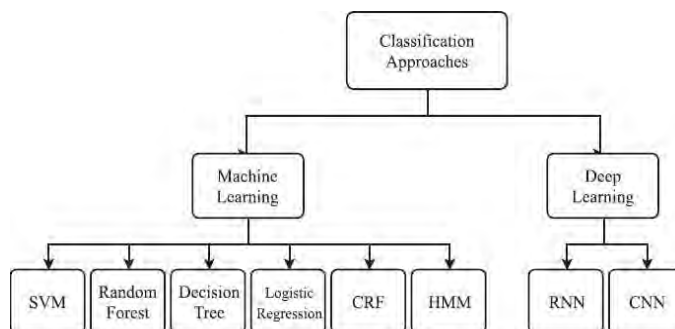


شکل ۲- توزیع چاپ مقالات در حوزه اخبار جعلی و تشخیص شایعات از سال ۲۰۱۱ تا ۲۰۱۸ (زوبیگا و همکاران، ۲۰۱۵)

بیشتر رویکردها با هدف شناسایی اخبار جعلی بر استفاده از ویژگی‌های محتوای طبقه‌بندی متمرکز شده‌اند. در حقیقت، طبق پژوهش ناپ (۱۹۴۴)، رویکردهای بسیار کمی برای کشف اخبار جعلی به مدل‌های صرفاً اجتماعی متکی بوده‌اند. در مقابل، رویکردهای شناسایی و تایید شایعه اغلب از ترکیبی از محتوا و ویژگی‌ها برای مدل‌های خود استفاده می‌کنند. واضح است، این امر به این دلیل است که جنبه اجتماعی شایعات ممکن است نقش اساسی در بهبود عملکرد تشخیص داشته باشد. غالباً محققان برای یافتن مناسب‌ترین مدل طبقه‌بندی اخبار جعلی برای مجموعه داده‌ها و کارها، مجبور به آزمایش الگوریتم‌های طبقه‌بندی مختلف هستند.

شکل ۳ رویکردهای مختلفی را که ما در تحلیل خود در نظر خواهیم گرفت، خلاصه می‌کند.

این رویکردها را می‌توان به رویکردهای طبقه‌بندی و سایر رویکردها دسته‌بندی کرد. رویکردهای طبقه‌بندی می‌توانند به نوبه خود مبتنی بر یادگیری ماشین و یادگیری عمیق باشند.



شکل ۳- رویکردهای موجود در ادبیات برای مسأله اخبار جعلی (زوبیگا و همکاران، ۲۰۱۵)

ثابت شده است که الگوریتم‌های یادگیری ماشین برای حل بسیاری از کارها در زمینه مهندسی اطلاعات می‌توانند مفید واقع شوند. از آنجایی که اولین رویکردها با تمرکز بر اعتبار در رسانه‌های اجتماعی (زوبیگا و همکاران، ۲۰۱۸) و کشف فریب با واسطه رایانه (زوبیگا و همکاران، ۲۰۱۸) نتایج امیدوارکننده‌ای را ارائه داده‌اند، تکنیک‌های یادگیری ماشین در تعدادی از تحقیقات در رابطه با مسئله کشف اطلاعات غلط مورد تست و ارزیابی قرار گرفتند. بیشتر رویکردهای یادگیری ماشین که برای اخبار جعلی و کشف شایعات به کار رفته‌اند، از یک استراتژی یادگیری نظارت شده استفاده می‌کنند. در جدول زیر به برخی از این الگوریتم‌ها و تلاش‌ها برای کشف اخبار جعلی می‌پردازیم.

جدول ۲- دسته‌بندی رویکردهای یادگیری ماشین سنتی برای اخبار جعلی

روش	منبع
Support Vector Machine (SVM)	Zubiaga & et al., 2018; Castillo, Mendoza & Poblete, 2011; Zhang & et al., 2012; Afroz, Brennan & Greenstadt, 2012; Kohavi & Quinlan, 2002
Decision Tree (DT)	Rubin & et al., 2016; Qin & et al., 2016; Zubiaga & et al., 2018; Yang & et al., 2012; Zhou & et al., 2003
Random Forest (FR)	Rubin & et al., 2016; Yang & et al., 2012; Chang & et al., 2016; Aker, Derczynski & Bontcheva, 2017; Ruchansky, Seo & Liu, 2017
Logistic Regression (LR)	Rubin & et al., 2016; Briscoe, Appling & Hayes, 2014; Zubiaga & et al, 2016 (b)
Conditional Random Field (CRF)	Aker, Derczynski & Bontcheva, 2017; Kwon & et al., 2013
Hidden Markov Models (HMMs)	Rubin, Chen & Conroy, 2015; Giasemidis & et al., 2016

۳-۱. رویکردهای مبتنی بر یادگیری عمیق

یادگیری عمیق^۱ یکی از مباحث تحقیقاتی است که به طور گسترده در زمینه یادگیری ماشین مورد بررسی قرار می‌گیرد. طبقه‌بندهای یادگیری عمیق به دلیل نتایج بسیار امیدوارکننده در تعدادی از زمینه‌های تحقیقاتی از جمله متن‌کاوی و پردازش زبان طبیعی^۲، در سال‌های اخیر به شدت مورد استفاده قرار گرفته‌اند. چارچوب‌های یادگیری عمیق به دلیل انتخاب ویژگی‌ها به صورت خودکار دقت بالاتری را نسبت به رویکردهای یادگیری سنتی بدست آورده‌اند. وظیفه استخراج ویژگی وقت‌گیر بوده و ممکن است منجر به ویژگی‌های مغرضانه شود (تاچینی^۳ و همکاران، ۲۰۱۷). این یک موضوع مهم برای کارهایی مانند اخبار جعلی و کشف شایعات است، جایی که شناسایی ویژگی‌های مربوط به آنالیز ممکن است یک چالش بزرگ‌تری را ایجاد کند. از سوی دیگر، چارچوب‌های یادگیری عمیق می‌توانند بازنمایی‌های پنهان را از ورودی‌های ساده‌تر بدست بیاورند (تاچینی و همکاران، ۲۰۱۷).

بنابراین، مسئله از مدل‌سازی ویژگی‌های ورودی مرتبط، به مدل‌سازی خود شبکه منتقل می‌شود که امکان حل کارآمد مسئله را فراهم می‌آورد. دو الگوریتم که در شبکه‌های عصبی مصنوعی مدرن بسیار مورد استفاده قرار می‌گیرند، شبکه‌های عصبی مکرر^۴ و شبکه‌های عصبی کانولوشنی^۵ هستند. جدول ۳ رویکردهای موجود در ادبیات می‌باشند که از شبکه‌های عصبی عمیق استفاده کرده‌اند:

جدول ۳- رویکردهای شبکه‌های عصبی عمیق

منبع	روش
Zeng, Starbird & Spiro, 2016; Zubiaga, Liakata & Procter, 2016 (a); Vosoughi, 2015; Ma & et al., 2016	Convolutional Neural Network (CNN)
Zhou & et al., 2003; Chang & et al., 2016; Tacchini & et al., 2017; Jacovi, Shalom & Goldberg, 2018; Kwon & et al., 2013; Aker, Derczynski & Bontcheva, 2017; Vosoughi, 2015; LeCun, Kavukcuoglu & Farabet, 2010	Recurrent Neural Networks (RNNs)

http://stim.gom.ac.ir

1. Deep Learning
2. Natural Language Processing
3. Tacchini
4. RNN
5. CNN

۴. نوآوری‌های پژوهش

- (۱) نمایش جدید متن در یک ساختار سه بعدی
 - (۲) اعمال عملیات کانولوشن در سطح جمله
 - (۳) در نظر گرفتن ویژگی‌های مکانی جمله برای طبقه‌بندی
 - (۴) در نظر گرفتن همجواری جملات برای استخراج ویژگی‌های اضافه.
- در واقع نوآوری اصلی پژوهش حاضر استفاده از ساختار سه بعدی بوده که در پژوهش‌های قبلی چنین ساختاری در نظر گرفته نشده است. این ساختار، مکان قرارگیری جملات در سند را نیز در نظر می‌گیرد که باعث طبقه‌بندی بهتر اسناد می‌شود. همچنین به دلیل تغییر در ورودی ساختار متن‌ها، ساختار شبکه‌های کانولوشنی از حالت یک بعدی به دو بعدی تغییر می‌کنند.

۵. سوالات پژوهش

- (۱) آیا در این پژوهش رویکردهای یادگیری ماشین، دقت و صحت بهتری در تشخیص اخبار جعلی فارسی نسبت به رویکردهای یادگیری عمیق دارند؟
- (۲) آیا استفاده از تنسور سه بعدی بر روی طبقه‌بندی متن‌های فارسی امکان‌پذیر بوده و موجب بهبود دقت و صحت تشخیص اخبار جعلی شده است؟
- (۳) آیا با تبدیل تنسور سه بعدی متون در شبکه‌های کانولوشن به تنسور سه بعدی، مشکل آنالیز در سطح جمله حل شده و باعث افزایش دقت و صحت تشخیص و طبقه‌بندی می‌شود؟

۶. روش‌شناسی تحقیق



شکل ۴- گردش کاری برای هر پروژه پردازش زبان طبیعی به صورت کلی

۶-۱. مجموعه داده‌های جمع‌آوری شده اخبار فارسی

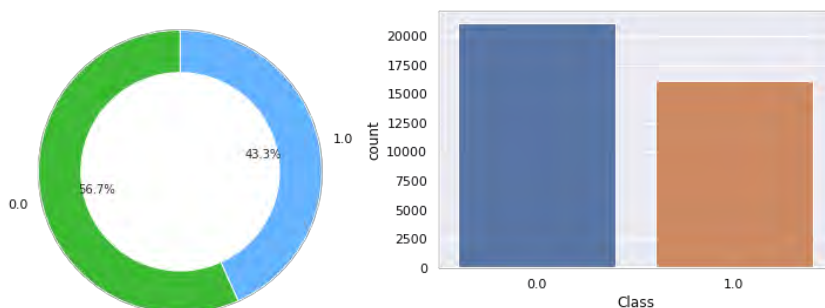
داده‌های جمع‌آوری شده در این پژوهش به مدت ۲ ماه از تاریخ ۱۳۹۹/۰۱/۰۱ لغایت ۱۳۹۹/۰۳/۰۵ به تعداد ۴۲۰۰۰ خبر فارسی بوده که از این تعداد، ۸۰ درصد برای آموزش و ۲۰ درصد برای تست مدل استفاده شده است. این داده‌ها در چندین مرحله جمع‌آوری و تکمیل شده است. در هر مرحله اطلاعات بیشتری به داده‌های اولیه اضافه شده است. داده‌ها در سطریهای مجزا

ذخیره گردیده که هر سطر را می‌توان متناظر با یک کاربر و یک داده‌ی ورودی برای آموزش مدل در نظر گرفت. در ادامه به ترتیب از جمع‌آوری داده‌های اولیه شروع کرده و سپس در هر مرحله تکمیل داده‌های اولیه و اضافه شدن ویژگی‌های مورد نیاز بررسی می‌شود و در نهایت اطلاعات آماری کل داده‌های جمع‌آوری شده مطرح خواهد شد.

تویتر اغلب به عنوان یک پلتفرم عمومی شناخته می‌شود و API‌های متفاوتی برای جمع‌آوری داده‌های آن معرفی شده است (اندروفر، ۲۰۱۷). در پژوهش حاضر از GetOldTweets و Tweepy برای جمع‌آوری داده‌ها استفاده گردید. بعضی از ابزارها دسترسی به توییت‌های قدیمی‌تر را فراهم می‌کنند. GetOldTweets یک ابزار کاملاً رایگان برای جمع‌آوری داده‌های توییت بوده که از خاصیت‌های جستجوی ترکیبی و جستجو براساس یک کلمه نیز پشتیبانی می‌کند و این اجازه را می‌دهد تا به توییت‌های طولانی مدت دست یابیم. این API اطلاعات بسیار مفیدی نظیر id (str), permalink (str), username (str), to (str), text (str), date (datetime) in UTC, retweets (int), favorites (int), mentions (str), hashtags (str), geo (str) را فراهم می‌آورد. ویژگی‌هایی که GetOldTweets استخراج می‌کند مفید، اما بسیار کم می‌باشند. برای این منظور از Tweepy برای استخراج برخی از ویژگی‌های مفید دیگر نظیر تعداد followerها و followingهای هر فرد استفاده شد. این ابزار قدرتمند نیز برای جمع‌آوری داده‌های توییت مورد استفاده قرار می‌گیرد که از مکانیزم OAuth برای احراز هویت استفاده می‌کند. برای جمع‌آوری داده‌ها چهار شهر مختلف تهران، اصفهان، قم و مشهد به عنوان منطقه مکانی گزینش گردید. همچنین دو هشتگ #کرونا و #کرونا_ویروس به عنوان هشتگ جمع‌آوری در نظر گرفته شد.

۲-۶. برچسب زدن

در این بخش داده‌ها، توسط گروهی از افراد خبره و با کمک سایت‌های خبری برچسب مناسب زده شده است. برچسب صفر به عنوان اخبار واقعی و برچسب یک به عنوان اخبار جعلی محسوب می‌گردد.



شکل ۴- آنالیز آماری مربوط به توییت‌های جمع‌آوری شده
 (در این دو نمودار + نشان‌دهنده کلاس Real و ۱ نشان‌دهنده کلاس Fake می‌باشد)

۳-۶. فاز پیش‌پردازش (استخراج متون و پیش‌پردازش)

در این گام پیش‌پردازش‌های رایج در حوزه پردازش متن نظیر حذف ایست واژه‌ها، حذف علائم نگارشی و توکن‌بندی بر روی متن‌های پرچسب‌دار اعمال شده است. معمولاً چند مرحله در زمینه پاک‌سازی و پیش‌پردازش داده‌های متنی وجود دارد. با این حال در این بخش نیز برخی از مهم‌ترین گام‌هایی را که در تحقیق پیش‌رو صورت گرفته به اختصار توضیح داده می‌شود. این گام‌ها به وفور در پروژه‌های NLP مورد بهره‌برداری قرار می‌گیرند. در این تحقیق از nltk و scikit-learn و spacy استفاده شده که جزء کتابخانه‌های تثبیت شده در حوزه NLP هستند.

۱-۳-۶. حذف کردن تگ‌های HTML

متن‌های ساخت‌نیافته غالباً شامل مقدار زیادی نویز هستند، به خصوص اگر از تکنیک‌هایی مانند اسکرپ کردن وب یا صفحه استفاده شود. تگ‌های HTML به طور معمول یکی از مؤلفه‌هایی هستند که ارزش زیادی در جهت درک و آنالیز متن اضافه نمی‌کنند.

۲-۳-۶. کاراکترهای ویژه

کاراکترهای ویژه و نمادها معمولاً کاراکترهای عددی- حرفی یا حتی در مواردی کاراکترهای عددی (بسته به مسئله) باعث افزایش نویز در متون می‌شوند که بایستی حذف گردد.

۳-۳-۶. حذف ایست‌واژه‌ها^۱

کلماتی که بی معنی هستند یا معنای خاصی ندارند، به خصوص وقتی ویژگی‌های معنایی از متن

استخراج می‌شود، به نام نامیده می‌شوند. این موارد معمولاً کلماتی هستند که فراوانی زیادی در متن دارند. به طور معمول این کلمات شامل حروف تعریف، ربط، اضافه و مواردی از این گروه هستند.

۴-۳-۶. نرمال‌سازی متن

با جمع‌بندی داده‌ها و مرتبط ساختن عملیات‌هایی همچون موارد مذکور، یک نرمال‌سازی متن برای پیش‌پردازش داده‌های متنی ایجاد می‌شود.

از دستور CountVectorizer که توسط کتابخانه learn scikit برای برداری‌سازی جملات فراهم شده، استفاده می‌گردد. این پیاده‌سازی کلمات هر جمله را دریافت کرده و یک واژه از همه کلمات یکتای موجود در جمله می‌سازد. این واژه را می‌توان برای ساخت یک بردار ویژگی از شمارش‌های کلمات، مورد استفاده قرار داد. CountVectorizer کار «توکن‌سازی»^۱ را انجام می‌دهد و به عبارت دیگر، جملات را به مجموعه‌ای از «توکن‌ها» جداسازی می‌کند. علاوه بر این، نشانه‌گذاری‌ها و کاراکترهای خاص را نیز حذف کرده و می‌تواند پیش‌پردازش را بر هر واژه اعمال کند.

۴-۶. جاسازی کلمات

یکی از مشکلات متن‌ها برای شبکه‌های عصبی این است که نمی‌توان آن‌ها را به صورت مستقیم به شبکه داد. برای این منظور آن‌ها را به صورت بردار تعبیه شده به شبکه می‌فرستند. راه‌های مختلف بردارسازی متن به صورت کلمات و کاراکترها بوده که در این پژوهش پس از بردارسازی با روش fasttext که نوعی جاسازی کلمات از پیش تعیین شده بوده و توسط شرکت فیسبوک ابداع شده، اقدام گردیده که فضای جاسازی^۲ نامیده می‌شود.

۵-۶. آموزش و تست مدل: تعریف مدل مبنا

هنگام انجام کار یادگیری ماشین، یک گام مهم، تعریف مدل مبنا است. مدل مبنا معمولاً شامل یک مدل ساده بوده که بعداً برای مقایسه با مدل پیشرفته‌تری که قصد آزمودن آن وجود دارد، مورد استفاده قرار می‌گیرد. در این شرایط، از مدل مبنا برای مقایسه آن با متدهای پیشرفته‌تری از جمله شبکه‌های عصبی (عمیق) بهره‌برداری می‌گردد. ابتدا، باید داده‌ها به مجموعه‌های «آموزش»^۳ و «آزمون»^۴ تقسیم شوند تا امکان ارزیابی صحت و قابلیت تعمیم مدل فراهم شود.

1. Tokenization
2. Embedding Space
3. Train
4. Test

۶-۶. ایجاد مدل یادگیری عمیق در شبکه کانولوشن

یک CNN دارای لایه‌های پنهانی است که «لایه‌های پیچشی»^۱ نامیده می‌شود. هنگامی که کاربر به تصاویر فکر می‌کند، کامپیوتر باید با ماتریس دو بعدی از اعداد کار کند. بنابراین، نیاز به راهی برای شناسایی ویژگی‌ها در این ماتریس است. این لایه‌های پیچشی قادر به شناسایی لبه‌ها، گوشه‌ها و دیگر انواع الگوها هستند که آن‌ها را به ابزاری ویژه مبدل می‌سازد. لایه پیچشی شامل چندین فیلتر است که در سراسر تصویر لغزش می‌کنند و قادر به شناسایی ویژگی‌های خاص هستند. این هسته اصلی روش و فرآیند ریاضی که در پیچش اتفاق می‌افتد، می‌باشد. با هر لایه پیچشی، شبکه قادر به شناسایی الگوهای پیچیده‌تر است. هنگام کار با داده‌های ترتیبی، مانند متن، با پیچش‌های یک‌بعدی کار می‌شود.

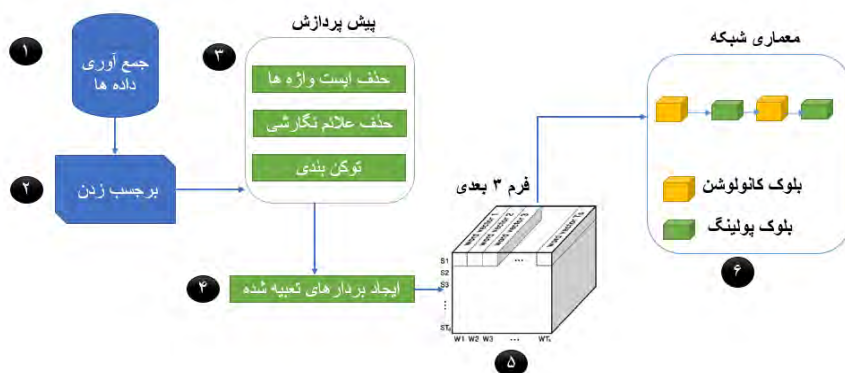
۶-۷. ساختار شبکه

در این ساختار، یک شبکه کانولوشنی بر روی ورودی‌های سه بعدی اعمال شد که برای این منظور نیاز بود در ساختار لایه کانولوشن و پولینگ، تغییراتی انجام شود که در بخش رویکرد پیشنهادی به این تغییرات پرداخته شده است.

- 1: Initialization weights to randomly generated value (small)
- 2: Set learning rate to a small value (positive)
- 3: Iteration $n = 1$; **Begin**
- 4: **for** $n < \text{max iteration OR Cost function criteria met}$, **do**
- 5: **for** image x_1 to x_i , **do**
- 6: a. Forward propagate through convolution, pooling and then fully connected layers
- 7: b. Derive Cost Function value for the image
- 8: c. Calculate error term $\delta^{(l)}$ with respect to weights for each type of layers.
- 9: Note that the error gets propagated from layer to layer in the following sequence
- 10: i. fully connected layer
- 11: ii. pooling layer
- 12: iii. convolution layer
- 13: d. Calculate gradient $\nabla_{w_k^{(l)}}$ and $\nabla_{b_j^{(l)}}$ for weights $\nabla_{w_k^{(l)}}$ and bias respectively for each layer
- 14: Gradient calculated in the following sequence
- 15: i. convolution layer
- 16: ii. pooling layer
- 17: iii. fully connected layer
- 18: e. Update weights
- 19: $w_{ji}^{(l)} \leftarrow w_{ji}^{(l)} + \Delta w_{ji}^{(l)}$
- 20: f. Update bias
- 21: $b_j^{(l)} \leftarrow b_j^{(l)} + \Delta b_j^{(l)}$

شکل ۵- شبه کد شبکه‌های کانولوشنی

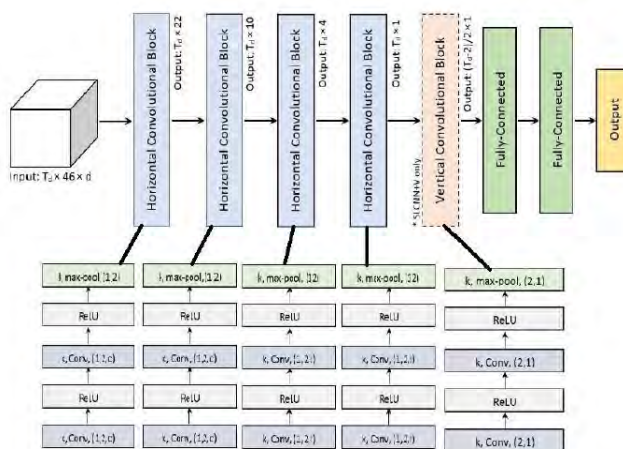
پس از توضیحات تکمیلی فوق، تمامی مراحل کار به صورت چارت نمایش داده شده (شکل چارچوب پیشنهادی) و به رویکرد پیشنهادی تنسور سه بعدی که ایده اصلی تحقیق است، پرداخته می‌شود.



شکل ۶- چارچوب روش پیشنهادی

۷. رویکرد پیشنهادی

مسئله تشخیص اخبار جعلی به عنوان یک مسئله طبقه‌بندی در پردازش زبان طبیعی شناخته می‌شود و هدف آن اختصاص یکی از کلاس‌های جعل و واقعی به متن‌های آزاد می‌باشد. شبکه‌های عصبی کانولوشنی به عنوان یکی از مهم‌ترین مدل‌های یادگیری عمیق، دقت بالایی را بر روی این مسائل بدست آورده است. در این شبکه‌ها کلمات به صورت کیسه‌ای از کلمات به مدل داده می‌شوند که هر کلمه با توجه به فضای برداری (Glove, Word2vec, fasttext) به ماتریس‌های دو بعدی تبدیل می‌شوند. یکی از مهم‌ترین چیزهایی که در این شبکه‌ها از بین می‌رود، مکان قرارگیری کلمات در سطح جملات است.

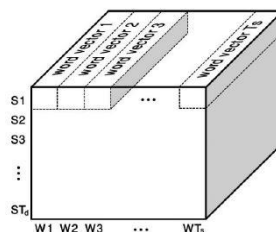


شکل ۷- الگوی اولیه از رویکرد پیشنهادی

نکته قابل توجه این است که شبکه‌های عصبی کانولوشنی انتخاب ویژگی‌های آن‌ها به صورت خودکار است. این شبکه‌ها در سطح کلمه کار می‌کنند و نمی‌توانند رابطه و فاصله بین جملات را در نظر بگیرند. برای این منظور نیاز است که این رابطه‌ها در نظر گرفته شود، چرا که ممکن است این رابطه‌ها دارای ویژگی‌های مفیدی باشند.

مدل پیشنهادی شبکه کانولوشنی مبتنی بر تئسور سه بعدی 3D-TCNN^۱ است. ایده کلیدی این رویکرد استفاده از اطلاعات مکانی هر جمله در اسناد می‌باشد. شایان ذکر است که در نظر گرفتن جملات در کنار هم می‌تواند اطلاعات اضافه‌ای را ایجاد کند که در مدل‌های مبتنی بر کیسه‌ای از کلمات، این اطلاعات قابل دسترس نیستند. پژوهش حاضر دو رویکرد مبتنی بر شبکه‌های عصبی کانولوشنی برای طبقه‌بندی اخبار جعل پیشنهاد می‌دهد. برای این منظور ابتدا به فرایند تولید تئسورهای سه بعدی که فرایند پردازش در سطح جمله را امکان‌پذیر می‌سازد، پرداخته می‌شود.

<http://stun.gom.ac.ir>



شکل ۸- تئسور سه بعدی

هر یک از اسناد ورودی پس از پیش‌پردازش به یک تنسور سه بعدی تبدیل می‌شود. شکل ۸ نمایش این تنسور می‌باشد. همان‌طور که در شکل ۸ نیز نشان داده شده است، در بعد اول سند، بعد دوم کلمات و بعد سوم بردارهای تعبیه شده قرار می‌گیرند.

یکی از محدودیت‌های شبکه‌های CNN این است که ابعاد ورودی‌ها باید با هم برابر باشد. از طرفی طول دنباله متن‌های ورودی با هم متفاوت است و برای این منظور دو حد آستانه در نظر گرفته می‌شود. حد آستانه اول که نشان‌دهنده تعداد جملات ورودی است و با Threshold Sentence (TS) نشان داده می‌شود و حد آستانه دوم که بیان‌کننده تعداد کلمات در یک جمله می‌باشد که با Threshold Word (TW) نشان داده خواهد شد.

پس از بررسی میزان داده‌های پرت^۱ طول اسناد به عنوان حد آستانه دوم در نظر گرفته می‌شود و میزان حد آستانه اول نیز از رابطه زیر بدست می‌آید:

$$T_s = \lfloor \mu + \sigma \rfloor \quad (۱)$$

در این رابطه μ میانگین تعداد جملات در کل اسناد و σ انحراف معیار می‌باشد. شایان ذکر است که برای انتخاب TS مقدار اوت لایر نیز می‌تواند در نظر گرفته شود.

معماری مدل پیشنهادی در شکل ۸ نشان داده شده است. در حالت کلی، در لایه ورودی، اسناد به فرم سه بعدی طبق الگوی بررسی شده در بخش قبل درمی‌آیند. چهار بلوک عمودی کانولوشنی که وظیفه استخراج ویژگی به ازای هر جمله واحد را بر عهده دارند، بر روی ورودی‌ها اعمال می‌شوند. این لایه‌ها وظیفه استخراج ویژگی‌ها به ازای هر یک از جملات ورودی را بر عهده دارند. همان‌طور که قبلاً گفته شد، در نظر گرفتن موقعیت مکانی کلمات در جمله‌ها می‌تواند اطلاعات مفیدی را برای طبقه‌بندی ایجاد کند. برای این منظور بلوک‌های افقی کانولوشنی بر روی خروجی بلوک‌های عمودی کانولوشنی اعمال می‌شود. در انتها دو لایه تماماً متصل جهت یادگیری ویژگی‌های استخراج شده بر روی خروجی‌های کانولوشن‌های افقی اعمال می‌شود. هر یک از بلوک‌های کانولوشنی افقی و عمودی به صورت جزئی دارای ۴ لایه و عملگر زیر می‌باشد:

لایه کانولوشن: این لایه وظیفه استخراج ویژگی‌های ضمنی از ماتریس ورودی و یا نقشه ویژگی‌های میانی است. عملگر کانولوشنی مد نظر در رویکرد پیشنهادی فیلترهایی به اندازه $w \in R^{s*t*d}$ را بر روی هر یک از نقشه ویژگی‌ها به اندازه $s*t$ اعمال می‌کند که برای این منظور از

رابطه زیر استفاده خواهد کرد:

$$x'_{i,j} = f(w \cdot x_{i,j:i+s-1,j+t-1} + b) \quad (2)$$

در این رابطه $x_{i,j:i+s-1,j+t-1}$ ویژگی‌های ادغام شده براساس یک بازه خاص، $b \in R$ بایاس و f یک تابع غیر خطی می‌باشد که برای این منظور از ReLU استفاده شده است. برای کانولوشن عمودی مقدار $s = 1$ و $t = 2$ و برای کانولوشن افقی مقدار $s = 2$ و $t = 1$ در نظر گرفته شده است. قابل توجه است که در کانولوشن عمودی در لایه اول d برابر اندازه بردارهای تعبیه شده و در سایر لایه‌ها برابر ۱ می‌باشد.

لایه pooling: یک لایه pooling معمولاً بعد از یک لایه کانولوشنی قرار می‌گیرد و از آن برای کاهش اندازه نقشه ویژگی feature map ها و پارامترهای شبکه استفاده می‌شود. استفاده از Max pooling و Average pooling در شبکه‌های کانولوشنی بسیار رایج است. در رویکرد پیشنهادی Max pooling با اندازه $pooling = 2$ بر روی آن جهت استخراج مهم‌ترین ویژگی‌ها از نقشه ویژگی‌های میانی اعمال می‌شود.

ویژگی‌های جدید از طریق روابط زیر ایجاد می‌شود:

$$\begin{cases} \max\{x_{i,2j-1}, x_{i,2j}\} \\ \max\{x_{2i-1,j}, x_{2i,j}\} \end{cases} \quad (3)$$

۷-۱. نتایج آزمایشات

در این بخش پس از معرفی محیط پیاده‌سازی، پارامترهای شبیه‌سازی معرفی خواهد شد. در ادامه نیز با معرفی معیارهای ارزیابی کارایی روش پیشنهادی براساس معیارهای ارزیابی موصوف بررسی و یافته‌ها با دیگر روش‌های مشابه مقایسه می‌گردد که برای این مقایسه‌ها از رویکرد یادگیری عمیق مبتنی بر شبکه کانولوشنی عمیق و دیگر رویکردهای یادگیری عمیق استفاده می‌شود.

۷-۲. محیط پیاده‌سازی

کراس یک کتابخانه سطح بالا، برای ایجاد ابزارهای شبکه‌های عصبی است که به زبان پایتون نوشته شده و قابلیت اجرا بر روی TensorFlow, CNTK و Theano را دارد. کراس شامل پیاده‌سازی‌های متعددی از بلوک‌های ساختار شبکه عصبی مانند: لایه‌ها، شیء‌ها، توابع فعال‌ساز، بهینه‌سازها و همچنین ابزارهای متعددی برای کار با تصاویر و داده‌های متنی است. این کتابخانه با نسخه‌های 2,7-3,6 پایتون سازگار می‌باشد و می‌تواند به صورت یکپارچه بر روی CPU ها و

GPUها اجرا شود. این API تمرکز خود را بر روی کاربرپسند بودن، ماژولاریتی و توسعه‌پذیری گذاشته است. کراس برای یادگیری و استفاده بسیار آسان می‌باشد، اما این سهولت استفاده باعث کاهش انعطاف‌پذیری آن نمی‌شود. کراس قابلیت ادغام با زبان‌های سطح پایین یادگیری عمیق به‌ویژه تنسورفلو را دارد. در پژوهش حاضر کراس برای ایجاد شبکه‌های عصبی استفاده شده است.

۷-۳. معیارهای ارزیابی

با در نظر داشتن پارامترهای زیر می‌توان معیارهای ارزیابی یک طبقه‌بند دودویی را به صورت زیر در نظر گرفت.

TP (True Positive): تعداد نمونه‌های مثبتی که مثبت پیش‌بینی شده‌اند.

FP (False positive): تعداد نمونه‌های منفی است که مثبت پیش‌بینی شده‌اند.

TN (True Negative): تعداد نمونه‌های منفی است که منفی دسته‌بندی شده‌اند.

FN (False Negative): تعداد نمونه‌های مثبتی است که منفی دسته‌بندی شده‌اند.

معیار precision: در یک طبقه‌بندی دودویی، precision کسری از اسناد پیش‌بینی شده مثبت که دقیقاً مثبت می‌باشند، نسبت به کل پیش‌بینی‌های مثبت انجام شده، می‌باشد (تاچینی، ۲۰۱۷).

$$Precision = \frac{TP}{TP+FP} \quad (4)$$

به طور کلی Precision بالا به معنی توانایی بالا برای پیش‌بینی درست است.

معیار Recall: نسبت تعداد اسناد پیش‌بینی شده مثبت به کل اسناد موجود شامل پیش‌بینی مثبت و منفی می‌باشد (تاچینی، ۲۰۱۷).

$$Recall = \frac{TP}{TP+FN} \quad (5)$$

Recall بالا به معنی این است که طبقه‌بند توانسته است میزان زیادی از اسناد را درست طبقه‌بندی کند.

F-Measure: میانگین هارمونیک نیز نامیده می‌شود که ترکیبی از دو معیار Precision و Recall می‌باشد. معمولاً به آن F-Score متعادل نیز می‌گویند (تاچینی، ۲۰۱۷).

$$F_{score} = 2 * \frac{Precision * Recall}{Precision + Recall} \quad (6)$$

Accuracy: یکی دیگر از معیارهایی است که در یک طبقه‌بند دودویی مورد استفاده قرار می‌گیرد. برای محاسبه‌ی آن نسبت TP و TN را به تعداد کل موارد پیش‌بینی شده در نظر می‌گیرند

(گارل^۱ و همکاران، ۲۰۱۸).

$$Accuracy = \frac{TP+FN}{TP+FP+FN+TN} \quad (V)$$

۸. بررسی نتایج حاصل از اجرای رویکرد پیشنهادی 3D-TCNN

جدول ۴ نشان‌دهنده دقت بدست آمده از رویکرد پیشنهادی اول بر روی مجموعه داده‌های مورد نظر می‌باشد. این رویکرد با ۵ رویکرد متفاوت دیگر که در پژوهش چو^۲ و همکاران (۲۰۱۴) معرفی شده‌اند، مورد مقایسه قرار گرفت. این رویکردها عبارتند از:

مدل Character base CNN: مدل دنباله‌ای از کاراکترهای رمزگذاری شده را به عنوان ورودی می‌پذیرند. رمزگذاری با تجویز الفبای اندازه m برای زبان ورودی انجام می‌شود و سپس با استفاده از رمزگذاری $m-1$ (یا رمزگذاری one-hot) هر کاراکتر را به مجموعه‌ای از بردارها تبدیل می‌کنیم. سپس، دنباله کاراکترها به دنباله‌ای از چنین بردارهایی با اندازه m با طول ثابت تبدیل می‌شود. الفبای مورد استفاده در تمام مدل‌ها شامل ۷۶ کاراکتر است که شامل ۳۲ حرف فارسی، ۱۰ رقم، ۳۳ کاراکتر دیگر که علائم نگارشی و غیره هستند، می‌باشند. این مدل جهت مقایسه با مدل‌های پیشنهادی مورد استفاده قرار گرفت. ساختار این مدل در پژوهش کاجکینا^۳، لیاکاتا^۴ و زوییاگا (۲۰۱۸) آمده است.

مدل Multi-Channel CNN: در پژوهش کیم^۵ (۲۰۱۴)، یک شبکه CNN ساده با یک لایه کانولوشن در بالای بردارهای کلمه بدست آمده از یک مدل زبانی غیرنظارتی آموزش داده شده است. این بردارها توسط میکلو^۶ و همکاران در ۱۰۰ میلیارد کلمه گوگل منتشر شده‌اند و در دسترس عموم قرار دارند. این معماری به عنوان یکی دیگر از معماری‌ها جهت مقایسه با رویکردهای پیشنهادی مورد استفاده قرار گرفت.

ماشین‌های بردار پشتیبان: یکی از متداول‌ترین روش‌های طبقه‌بندی در برخی از زمینه‌های تحقیقاتی است. SVM ها طبقه‌بندهای تبعیض‌آمیز هستند که به طور رسمی توسط یک مرز جداکننده تعریف می‌شوند و سعی در یافتن مرزی دارند که بیشترین حاشیه امن را بین مجموعه داده‌ها داشته باشد.

1. Gorrell
2. Cho
3. Kochkina
4. Liakata
5. Kim
6. Mikolov

درخت تصمیم‌گیری: یکی دیگر از خانواده‌های الگوریتم‌هایی که به طور گسترده، به ویژه برای کارهای تجزیه و تحلیل شایعه مورد مطالعه قرار گرفته‌اند، درخت تصمیم‌گیری است (ژاو و همکاران، ۲۰۰۳). درخت تصمیم به منظور تعیین کلاس، تقسیم بازگشتی روی مقادیر ویژگی‌ها انجام می‌دهد. درخت تصمیم‌گیری برای داده‌هایی توسط الگوریتم‌هایی مانند J48 (C4.5) (زوبیگا و همکاران، ۲۰۱۶) تولید می‌شود. علی‌رغم سادگی نسبی آنها نسبت به سایر طرح‌های یادگیری ماشین، آن‌ها عملکرد رقابتی را در بسیاری از وظایف بدست آورده‌اند.

جنگل تصادفی: جنگل‌های تصادفی (ژاو و همکاران، ۲۰۰۳) در تعدادی از کارهای تحلیل شایعه به کار گرفته شده است. یک جنگل تصادفی مجموعه‌ای از چندین درخت تصمیم است که با ترکیب پیش‌بینی‌ها از هر درخت، خروجی نهایی را بدست می‌آورد. مطالعات تطبیقی در مورد الگوریتم‌های مختلف یادگیری ماشین برای شایعات و کارهای خبری جعلی نشان می‌دهد که جنگل‌های تصادفی می‌تواند به عنوان یک طبقه‌بند قوی برای تشخیص مورد استفاده قرار گیرد. رویکرد پیشنهادی به دلیل در نظر گرفتن موقعیت جمله‌ها در متن، در سه معیار ارزیابی بررسی شده دقت بالاتری را بدست آورده است.

جدول ۴- نتایج رویکرد پیشنهادی

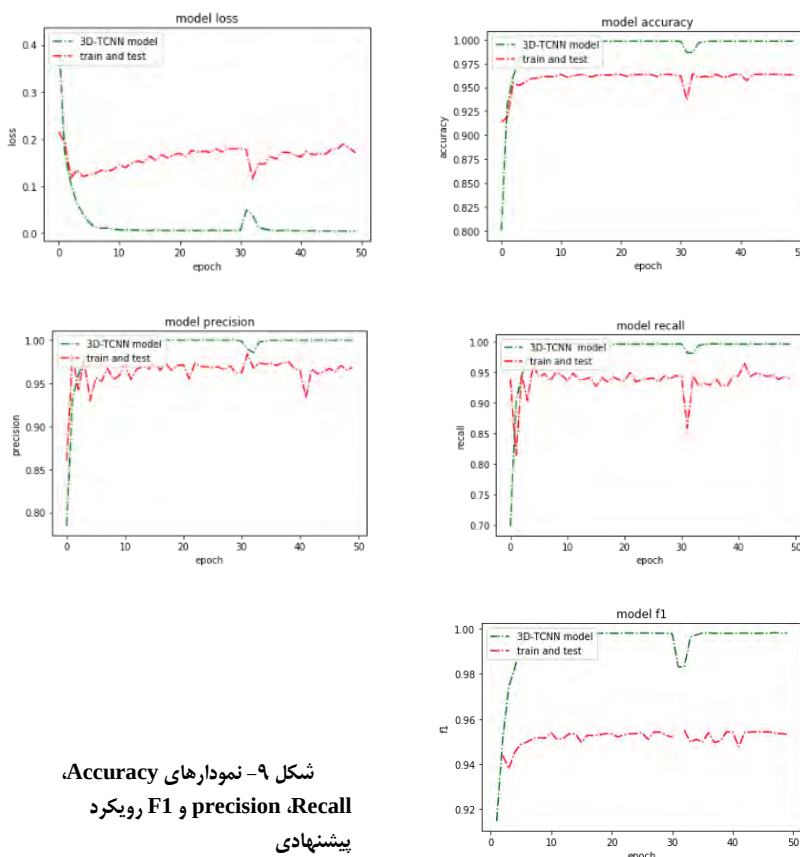
Methods	Accuracy	Precision	Recall	F1
Decision Tree	0.9200	0.9200	0.9200	0.9200
Random Forest	0.9100	0.9100	0.9100	0.9100
SVM	0.9100	0.9300	0.9100	0.9200
Char based model	0.8400	0.9300	0.7700	0.8700
Multi-channel CNN	0.8800	0.8800	0.8800	0.8800
3D-TCNN	0.9476	0.9430	0.9362	0.9388

در پاسخ به سوال ۱ تحقیق:

طبق جدول ۴، سه رویکرد یادگیری ماشین (Decision Tree, Random Forest و SVM) طبقه‌بندی دقت تقریباً برابری را بدست آورده‌اند. این در حالی است که رویکردهای مبتنی بر شبکه‌های کانولوشنی معمولی مانند Character based model و Multi-Channel CNN دقت کم‌تری نسبت به آن‌ها داشته‌اند. این میزان اختلاف ناشی از عدم استخراج ویژگی‌های مناسب برای رویکردهای یادگیری عمیق کانولوشنی می‌باشد. در واقع این رویکردها نتوانسته‌اند ویژگی‌های مناسب را استخراج کنند. از طرف دیگر، این اختلاف نشان‌دهنده قدرت رویکردهای یادگیری ماشین

می‌باشد که می‌توانند بر روی مسائل طبقه‌بندی به دقت مناسبی برسند.

در پاسخ به سؤال ۲ و ۳ پژوهش، رویکرد پیشنهادی 3D Tensor با توجه به در نظر گرفتن مکان قرارگیری جملات توانسته نسبت به سایر رویکردها به دقت بهتری دست یابد. تغییرات انجام شده در لایه‌های شبکه‌های کانولوشنی باعث شده است شبکه‌ی کانولوشنی بتواند ویژگی‌های مناسبی را استخراج کند و نسبت به رویکردهای شبکه‌های کانولوشنی معمولی و رویکردهای یادگیری ماشین معمولی به دقت بهتری دست یابد. یکی از مهم‌ترین پارامترها و معیارها برای ارزیابی نتایج مقایسه F1 می‌باشد که رویکرد پیشنهادی توانست به ۰/۹۴۶۱ برسد و نسبت به بهترین رویکرد موجود که ۰/۹۲ می‌باشد، مقدار قابل توجهی بهبود پیدا کرده است. نمودار شماره ۹ مربوط به میزان خطای مدل پیشنهادی برای ۵۰ مرحله اجرا است. همچنین نمودارهای مربوط به Recall، Accuracy، precision و F1 رویکرد پیشنهادی در ادامه آورده شده است.



شکل ۹- نمودارهای Accuracy، Recall، precision و F1 رویکرد پیشنهادی

با توجه به نمودارهای بدست آمده از رویکرد پیشنهادی می‌توان نتیجه گرفت که رویکرد پیشنهادی توانسته به خوبی مسأله پیش‌پرازش^۱ را کنترل کند. در ایپاک-های بالاتر از ۵۰ امکان پیش‌پرازش وجود دارد که برای این منظور می‌توان با تنظیم میزان Dropout این مسأله را تا حدودی حل کرد.

۹. نتیجه‌گیری

در این پژوهش رویکردی مبتنی بر تنسورهای سه بعدی برای طبقه‌بندی اخبار جعلی مورد بررسی قرار گرفت. همچنین به نتایج حاصل از دو رویکرد پیشنهادی 3D-TCNN بر روی مجموعه داده‌های مورد نظر با توجه به معیارهای ارزیابی پرداخته شد. نتایج نشان می‌دهد که در حالت کلی رویکرد پیشنهادی عملکرد بهتری نسبت تمامی رویکردهای پیشنهادی قبلی از جمله الگوریتم‌های پایه و الگوریتم‌های یادگیری ماشین داشته و به بیش از ۹۴ درصد دقت و صحت دست یافته است. همچنین با توجه به اینکه داده‌های مربوط به اخبار فارسی کرونا بسیار کم بوده، طی بررسی‌ها مشخص شد که کار یادگیری عمیق در مورد تشخیص اخبار جعلی فارسی به شدت کم بوده و شاید نادر می‌باشد و در مورد اخبار انگلیسی بیشتر از رویکردهای ترکیبی بین دو الگوریتم یادگیری عمیق استفاده شده و بسیار کم تغییر در ساختار الگوریتم از جمله افزایش یا کاهش پارامترها و سایر موارد، بررسی شده است. با این حال طی بررسی مقالات مشابه در زمینه تشخیص اخبار جعلی نتایج با دقت و صحت حدود ۹۳ درصد بدست آمده است. این رویکرد با پارامترهای متفاوتی مورد ارزیابی و تست قرار گرفت که نتایج موجود در جداول بررسی شده بهترین مقدار بدست آمده با این پارامترها بوده است. همچنین یکی از مهم‌ترین مشکلات رویکرد پیشنهادی، وجود لایه‌های کاملاً متصل می‌باشد که می‌تواند حجم محاسبات را به شدت افزایش دهد. برای این منظور می‌توان با حذف این لایه‌ها و جایگذاری رویکردهایی مانند ماشین بردار پشتیبان، به دقت بهتری دست یافت و این می‌تواند در پژوهش‌های بعدی بر روی رویکرد پیشنهادی اعمال شود.

<http://stjm.gom.ac.ir>

References

- Afroz, S.; Brennan, M. & Greenstadt, R. (2012). **Detecting hoaxes, frauds, and deception in writing style online**. In: 2012 IEEE Symposium on Security and Privacy (pp. 461-475). IEEE.
- Aker, A.; Derczynski, L. & Bontcheva, K. (2017). **Simple open stance classification for rumour analysis**. Proceedings of the International Conference Recent Advances in Natural Language Processing: 31-39. **DOI:** 10.26615/978-954-452-049-6_005.
- Allcott, H. & Gentzkow, M. (2017). Social media and fake news in the 2016 election. *Journal of economic perspectives*, 31(2): 211-36.
- Allport, G.W. & Postman, L. (1946). An analysis of rumor. *Public opinion quarterly*, 10(4): 501-517.
- Andorfer, A (2017). Spreading like wildfire: Solutions for abating the fake news problem on social media via technology controls and government regulation. *Hastings LJ*, 69: 1409.
- Berkowitz, D. & Schwartz, D.A. (2016). Miley, CNN and The Onion: When fake news becomes realer than real. *Journalism practice*, 10(1): 1-17.
- Briscoe, E.J.; Appling, D.S. & Hayes, H. (2014). **Cues to deception in social media communications**. In: 2014 47th Hawaii international conference on system sciences (pp.1435-1443). IEEE.
- Castillo, C.; Mendoza, M. & Poblete, B. (2011). **Information credibility on twitter**. In: Proceedings of the 20th international conference on World wide web (pp.675-684).
- Chang, C.; Zhang, Y.; Szabo, C. & Sheng, Q.Z. (2016). **Extreme user and political rumor detection on twitter**. In: International Conference on Advanced Data Mining and Applications (pp.751-763). Springer, Cham.
- Cho, K.; Van Merriënboer, B.; Gulcehre, C.; Bahdanau, D.; Bougares, F.; Schwenk, H. & Bengio, Y. (2014). **Learning phrase representations using RNN encoder-decoder for statistical machine translation**. Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP). 1724-1734. **DOI:** 10.3115/v1/D14-1179.
- Conroy, N.K.; Rubin, V.L. & Chen, Y. (2015). Automatic deception detection: Methods for finding fake news. *Proceedings of the Association for Information Science and Technology*, 52(1): 1-4.
- Derczynski, L. & Bontcheva, K. (2014). **Pheme: Veracity in Digital Social Networks**. In: UMAP workshops.
- Giasemidis, G.; Singleton, C.; Agrafiotis, I.; Nurse, J.R.; Pilgrim, A.; Willis, C. & Greetham, D.V. (2016). **Determining the veracity of rumours on Twitter**. In: International Conference on Social Informatics (pp.185-205). Springer, Cham.
- Gorrell, G.; Bontcheva, K.; Derczynski, L.; Kochkina, E.; Liakata, M. & Zubiaga, A. (2018). **Rumoureval 2019: Determining rumour veracity and support for rumours**. Proceedings of the 13th International Workshop on Semantic Evaluation. 845-854. **DOI:** 10.18653/v1/S19-2147.
- Jacovi, A.; Shalom, O.S. & Goldberg, Y. (2018). **Understanding convolutional neural networks for text classification**. Proceedings of the 2018 EMNLP Workshop BlackboxNLP: Analyzing and Interpreting Neural Networks for NLP. pp: 56-65. **DOI:** 10.18653/v1/W18-5408.
- Kim, Y. (2014). **Convolutional neural networks for sentence classification**. CoRR abs/1408.5882 (2014). Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP). 1746-1751. **DOI:** 10.3115/v1/D14-1181.

- Knapp, R.H. (1944). A psychology of rumor. *Public opinion quarterly*, 8(1): 22-37.
- Kochkina, E.; Liakata, M. & Zubiaga, A. (2018). **All-in-one: Multi-task learning for rumour verification**. Proceedings of the 27th International Conference on Computational Linguistics: 3402-3413.
- Kohavi, R. & Quinlan, J.R. (2002). **Data mining tasks and methods: Classification: decision-tree discovery**. In: Handbook of data mining and knowledge discovery: 267-276.
- Kshetri, N. & Voas, J. (2017). The economics of "fake news". *IT Professional*, 19(6): 8-12.
- Kucharski, A. (2016). Study epidemiology of fake news. *Nature*, 540(7634): 525-525.
- Kwon, S.; Cha, M.; Jung, K.; Chen, W. & Wang, Y. (2013). **Prominent features of rumor propagation in online social media**. In: 2013 IEEE 13th international conference on data mining (pp. 1103-1108). IEEE.
- LeCun, Y.; Kavukcuoglu, K. & Farabet, C. (2010). **Convolutional networks and applications in vision**. In: Proceedings of 2010 IEEE international symposium on circuits and systems (pp.253-256). IEEE.
- Ma, J.; Gao, W.; Mitra, P.; Kwon, S.; Jansen, B.J.; Wong, K.F. & Cha, M. (2016). **Detecting rumors from microblogs with recurrent neural networks**. Proceedings of the Twenty-Fifth International Joint Conference on Artificial Intelligence (IJCAI-16): 3818-3824.
- Meyer, J.K. (1969). **Bibliography on the urban crisis: The behavioral, psychological, and sociological aspects of the urban crisis**. In: proceedings Meyer 1969 Bibliography OT.
- Pogue, D. (2017). How to Stamp Out Fake News. *Scientific American*, 316(2): 24-24.
- Qin, Y.; Wurzer, D.; Lavrenko, V. & Tang, C. (2016). **Spotting rumors via novelty detection**. Vol. abs/1611.06322. n. pag.
- Rapoza, K. (2017). **Can 'fake news' impact the stock market?**. Available at: <https://www.forbes.com/sites/kenrapoza/2017/02/26/can-fake-news-impact-the-stock-market/?sh=293d89802fac>.
- Rubin, V.L.; Chen, Y. & Conroy, N.K. (2015). Deception detection for news: three types of fakes. *Proceedings of the Association for Information Science and Technology*, 52(1): 1-4.
- Rubin, V.L.; Conroy, N.; Chen, Y. & Cornwell, S. (2016). **Fake news or truth? using satirical cues to detect potentially misleading news**. In: Proceedings of the second workshop on computational approaches to deception detection: 7-17.
- Ruchansky, N.; Seo, S. & Liu, Y. (2017). **Csi: A hybrid deep model for fake news detection**. In: Proceedings of the 2017 ACM on Conference on Information and Knowledge Management: 797-806.
- Sharma, K.; Qian, F.; Jiang, H.; Ruchansky, N.; Zhang, M. & Liu, Y. (2019). Combating fake news: A survey on identification and mitigation techniques. *ACM Transactions on Intelligent Systems and Technology (TIST)*, 10(3): 1-42.
- Shu, K.; Sliva, A.; Wang, S.; Tang, J. & Liu, H. (2017). Fake news detection on social media: A data mining perspective. *ACM SIGKDD explorations newsletter*, 19(1): 22-36.
- Siering, M.; Koch, J.A. & Deokar, A.V. (2016). Detecting fraudulent behavior on crowdfunding platforms: The role of linguistic and content-based cues in static and dynamic contexts. *Journal*

- of Management Information Systems, 33(2): 421-455.
- Tacchini, E.; Ballarin, G.; Della Vedova, M.L.; Moret, S. & de Alfaro, L. (2017). **Some like it hoax: Automated fake news detection in social networks**. In: Proceeding in Conference SoGood 2017 - Second Workshop on Data Science for Social Good, Vol. 1960.
- Vosoughi, S. (2015). **Automatic detection and verification of rumors on Twitter**. Doctoral dissertation. Massachusetts Institute of Technology.
- Vosoughi, S.; Roy, D. & Aral, S. (2018). The spread of true and false news online. *Science*, 359 (6380): 1146-1151.
- Waldrop, M.M. (2017). News Feature: The genuine problem of fake news. *Proceedings of the National Academy of Sciences*, 114(48): 12631-12634.
- Yang, F.; Liu, Y.; Yu, X. & Yang, M. (2012). **Automatic detection of rumor on sina weibo**. In: Proceedings of the ACM SIGKDD workshop on mining data semantics: 1-7.
- Zeng, L.; Starbird, K. & Spiro, E. (2016). **# unconfirmed: Classifying rumor stance in crisis-related social media messages**. *Proceedings of the International AAAI Conference on Web and Social Media*, 10(1). No.1.
- Zhang, H.; Fan, Z.; Zheng, J. & Liu, Q. (2012). An improving deception detection method in computer-mediated communication. *Journal of Networks*, 7(11): 1811.
- Zhou, L.; Twitchell, D.P.; Qin, T.; Burgoon, J.K. & Nunamaker, J.F. (2003). **An exploratory study into deception detection in text-based computer-mediated communication**. In: 36th Annual Hawaii International Conference on System Sciences. IEEE.
- Zhou, X. & Zafarani, R. (2018). Fake news: A survey of research, detection methods, and opportunities. *ACM Computing Surveys*, 53(109): 1-40. DOI: 10.1145/3395046.
- Zubiaga, A.; Aker, A.; Bontcheva, K.; Liakata, M. & Procter, R. (2018). Detection and resolution of rumours in social media: A survey. *ACM Computing Surveys (CSUR)*, 51(2): 1-36.
- Zubiaga, A.; Liakata, M. & Procter, R. (2016). Learning reporting dynamics during breaking news for rumour detection in social media. *ACM Transactions on Management Information Systems*, 12(8): 1-16. DOI: 10.1145/3416703 (a).
- Zubiaga, A.; Liakata, M.; Procter, R.; Bontcheva, K. & Tolmie, P. (2015). **Towards detecting rumours in social media**. In: Workshops at the Twenty-Ninth AAAI conference on artificial intelligence.
- Zubiaga, A.; Liakata, M.; Procter, R.; Wong Sak Hoi, G. & Tolmie, P. (2016). Analysing how people orient to and spread rumours in social media by looking at conversational threads. *PLoS one*, 11(3): e0150989 (b).

http://stim.gom.ac.ir

استناد به این مقاله

DOI: 10.22091/stim.2021.7014.1592

متقی، وحید؛ اسماعیلی، مهدی؛ بازایی، قاسمعلی؛ افشار کاظمی، محمدعلی (۱۴۰۰). ارائه رویکرد تسور سه بعدی برای طبقه‌بندی و تشخیص اخبار جعلی: مطالعه موردی اخبار فارسی در حوزه کرونا ویروس. *علوم و فنون مدیریت اطلاعات*، ۷(۴): ۲۲۱-۲۵۰.