

Evaluating Trust Management in Social Networks Using Eviden Theory

Javad Hamidzadeh^{*1}, Abdolreza Abdolsalami², Amir Dahaki toroghi³, Mona Zendeheh⁴

¹ Faculty of Computer Engineering and Information Technology, Sadjad University, Mashhad, Iran.

² Faculty of Computer Engineering and Information Technology, Sadjad University, Mashhad, Iran.

³ Faculty of Computer Engineering and Information Technology, Sadjad University, Mashhad, Iran.

⁴ Faculty of Computer Engineering and Information Technology, Sadjad University, Mashhad, Iran.

* Corresponding author

Short Abstract

Various methods have been presented to evaluate trust management between users in social networks. Individual and subjective diagnosis is less considered in the evaluation of trust and often a general and general model is presented for all users. trust evaluation without considering the personal and mental characteristics of users is not effective. in the proposed method of this research, users' characteristics are calculated and their importance is determined using fuzzy rough set theory. users' characteristics are combined and aggregated by considering their importance and using Dempster Shafer theory of evidence. the sets of trust and distrust values and ambiguity are evaluated to determine the degree of trust. the final values are used to make a trust decision and create a secure connection. the level of user trust in the entire network is determined by all users. the comprehensive performance of the proposed method and RTARS, ABC, DSL-STM and AUTOMATA algorithms have been compared in four evaluation indices and in 10 independent implementations. the obtained results show the improvement of the trust decision and the creation of safe communication in social networks. the proposed algorithm has been able to correctly decide the trust of users in social networks with 92.54% accuracy. the experimental results show that the proposed method is able to infer trust more accurately than the previous methods.

Keywords

Trust management, dempster shafer evidence theory, user characteristics, social networks.

1- Short Introduction

Today, due to the online communication and the existence of the internet platform and its use in all financial centers and commercial companies, the issue of the security of such centers is sensitive and important. secure communication and information protection are one of the main challenges of these centers, so that a lot of funds are spent by these companies to create a secure platform for communication, while sometimes there are many weaknesses and wrong trust in the offending users, causing a lot of losses. has been to these centers. meanwhile, creating a safe platform and increasing the trust factor is effective, and the scope of this research article can be used in social networks and professional communication software, and in a practical way, it can improve the security and communication factor of these software and websites.

2- Proposed Work and Methodology

To evaluate the trust in the proposed method, the characteristics of the users are used, and the evaluation of the trust is done based on the characteristics, reactions, and individual behavior of the users. In this model, first, the computable features of users are measured and the range of features is determined. The feature set includes activity level, follower, group, and reputation. By calculating the similarity of the feature values, the trust evidence is generated. Considering that the criteria of trust are qualitative and the decision on qualitative criteria lacks the necessary accuracy, in this article, we have been able to improve the accuracy in trust decision-making up to 92.54% by using the calculation of individual characteristics and fuzzy rough set theory to weight the characteristics and combine them using evidence theory.

3- Conclusion

In this article, a new method for evaluating trust is presented by considering the individual and mental conditions of users instead of a general evaluation for all users. The proposed method has been able to correctly decide whether to trust or not trust the user with an accuracy of 92.54. The proposed method was compared and evaluated with 4 other models in trust evaluation (RTARS model, ABC model, DSL-STM model, and A-automat model). The results show the superiority of the proposed method over other methods in terms of precision, accuracy, and average error criteria.

4- References

- [18]. S.Ahmadian, M.Afsharchi, M.Afsharchi, "an effective social recommendation method based on user reputation model and rating profile enhancement." *Journal of Information Science*, pp. 1–36, 2018.
- [24]. Sh.Saeidi, "A new model for calculating the maximum trust in Online Social Networks and solving by Artificial Bee Colony algorithm." *Computational Social Networks*, Vol 7, No, 3, 2020.
- [1]. N.Fatehi, H.S.Shahhoseini, J.Wei, C.T.Chang, "An automata algorithm for generating trusted graphs in online social networks." *Applied Soft Computing* 118 : 108475, 2022.
- [15]. W.Abdelghani, I.Amous, C.A.Zayani, F.Sèdes, G.Roman-Jimenez, "Dynamic and scalable multi-level trust management model for Social Internet of Things." *The Journal of Supercomputing*, 78(6), 8137-8193, 2022.

ارزیابی مدیریت اعتماد در شبکه‌های اجتماعی با استفاده از نظریه شواهد

جواد حمیدزاده

دانشیار، دانشکده مهندسی کامپیوتر و فناوری اطلاعات، دانشگاه سجاد، مشهد، ایران.

عبدالرضا عبدالسلامی

دانشجوی ارشد مهندسی کامپیوتر گرایش رایانش امن، دانشگاه سجاد، مشهد، ایران.

امیر دهکی طرقی

دانشجوی ارشد مهندسی کامپیوتر گرایش رایانش امن، دانشگاه سجاد، مشهد، ایران.

منا زنده دل

دانشجوی ارشد مهندسی کامپیوتر گرایش هوش مصنوعی و رباتیک، دانشگاه سجاد، مشهد، ایران.

چکیده

با ظهور شبکه‌های ارتباطی و اتصال گسترده رایانه‌ها و وسایل همراه، همواره امنیت کاربران و امنیت ارتباط آنان موردتوجه است. در حوزه شبکه‌های اجتماعی، اعتماد کاربران از اصلی‌ترین مسائل ارتباط کاربران است و مدیریت اعتماد، نقش اساسی در این زمینه ایفا می‌کند. در حال حاضر مدیریت اعتماد از مباحث تعیین کننده در حوزه امنیت نرم و ارتباطات امن است. در امنیت نرم، اعتماد و کسب باور بر اساس ویژگی‌ها و رفتار کاربران شکل می‌گیرد. روش‌های مختلفی برای ارزیابی مدیریت اعتماد بین کاربران در شبکه‌های اجتماعی ارائه شده است. تشخیص فردی و ذهنی در ارزیابی اعتماد کمتر موردتوجه قرار گرفته و اغلب مدلی کلی و عمومی برای همه کاربران ارائه شده است. ارزیابی اعتماد بدون در نظر گرفتن خصوصیت‌های فردی و ذهنی کاربران کارایی لازم را ندارد. در روش پیشنهادی این پژوهش، ویژگی‌های کاربران محاسبه می‌شود و با استفاده از نظریه مجموعه خشن فازی میزان اهمیت آنها تعیین می‌گردد. ویژگی‌های کاربران با در نظر گرفتن میزان اهمیت آنها و با استفاده از نظریه شواهد دمپستر شفر، ترکیب و تجمیع می‌شوند. مجموعه‌های مقادیر اعتماد و بی‌اعتمادی و ابهام برای تعیین درجه اعتماد ارزیابی می‌شوند. مقادیر نهایی به‌منظور تصمیم اعتماد و ایجاد ارتباط امن مورد استفاده قرار می‌گیرند. میزان اعتماد کاربر در کل شبکه توسط همه کاربران تعیین می‌گردد. عملکرد جامع روش پیشنهادی و الگوریتم‌های RTARS، ABC، DSL-STM و AUTOMATA در چهار شاخص ارزیابی و در ۱۰ اجرای مستقل مقایسه شده‌اند. نتایج به‌دست‌آمده نشان‌دهنده بهبود تصمیم اعتماد و ایجاد ارتباط امن در شبکه‌های اجتماعی است. الگوریتم پیشنهادی توانسته است با دقت ۹۲٫۵۴٪ در مورد اعتماد کاربران در شبکه‌های اجتماعی به‌درستی تصمیم بگیرد. نتایج تجربی نشان‌دهنده آن است که روش پیشنهادی قادر به استنتاج اعتماد با دقت بالاتری نسبت به روش‌های قبلی است.

کلمات کلیدی

مدیریت اعتماد، نظریه شواهد دمپستر شفر، مجموعه خشن فازی، ویژگی کاربر، شبکه‌های اجتماعی.

نام نویسنده مسئول: جواد حمیدزاده

ایمیل نویسنده مسئول: j_hamidzadeh@sadjad.ac.ir

تاریخ ارسال مقاله: ۱۴۰۳/۰۳/۰۸

تاریخ(های) اصلاح مقاله: ۱۴۰۲/۰۶/۱۱

تاریخ پذیرش مقاله: ۱۴۰۲/۰۷/۲۲

۱- مقدمه

کاربران به‌مرور زمان با داده‌ها و اطلاعاتشان می‌توانند موقعیت هر کاربر را نشان دهد.

پیش‌بینی رابطه اعتماد در شبکه‌های اجتماعی برخط موردتوجه زیادی قرار گرفته است؛ زیرا به کاربران اجازه می‌دهد از تجزیه و تحلیل پیچیده در تصمیم‌گیری فرار کنند [۲]. به لحاظ مفهومی، اعتماد بیانگر یک نوع سرمایه ارتباطی بین کاربران است. در حوزه شبکه‌های اجتماعی، مدیریت اعتماد و تعیین درجه‌ی اعتماد از مباحث اصلی است. میزان اعتماد کاربران به یکدیگر از طریق مدیریت اعتماد و ارزیابی آن محاسبه می‌شود. ارزیابی اعتماد به‌عنوان

با ظهور شبکه‌های ارتباطی و اتصال گسترده رایانه‌ها و وسایل همراه، همواره امنیت کاربران و ارتباط آنان موردتوجه بوده است. شبکه‌های اجتماعی بر اساس انواع مختلفی از وابستگی و علاقه‌مندی شکل گرفته‌اند. شبکه‌های اجتماعی شکل مجازی‌سازی شده از ساختار اجتماعی فیزیکی دنیای واقعی است [۱]. باتوجه به توسعه شبکه‌های اجتماعی، رفتار کاربران و دسترسی آنان اهمیت زیادی پیدا می‌کند و این شبکه‌ها قادرند گروه‌های مختلف با ایده‌ها و رفتارهای متشابه تولید کنند و یک جمعیت مجازی همگن تشکیل دهند. تحلیل رفتار

را نادیده گرفته‌اند و برای سهولت کار و ارزیابی، قطعیت را مورد وثوق قرار نمی‌دهند، این درحالی‌که ممکن است ابهامات و عدم قطعیت و ارزیابی ناقص در تصمیم اعتماد تعیین‌کننده و مؤثر باشد.

مشارکتهای اصلی این مقاله را میتوان به شرح زیر خلاصه کرد:

- استفاده از ویژگی‌های فردی کاربران برای محاسبه اعتماد
 - استفاده از مجموعه خشن فازی برای محاسبه میزان تاثیر هر ویژگی
 - تولید شواهد اعتماد بر اساس محاسبه شباهت مقادیر ویژگی‌ها
 - ترکیب شواهد اعتماد با استفاده از نظریه شواهد دمپستر شفر
- در این پژوهش ما توانسته‌ایم با استفاده از محاسبه ویژگی‌های فردی و نظریه مجموعه خشن فازی برای وزن‌دهی ویژگی‌ها و ترکیب آن‌ها با استفاده از نظریه شواهد دقت را در تصمیم‌گیری اعتماد تا ۹۲/۵۴ درصد بهبود دهیم. ساختار ادامه مقاله به شرح ذیل است: در بخش دوم، به‌مرور کارهای مرتبط مورد بررسی قرار گرفته است. در بخش سوم، مفاهیم پایه حوزه تحقیق تعریف می‌شوند. بخش چهارم، راهکار پیشنهادی ارزیابی اعتماد مبتنی بر نظریه شواهد بیان می‌شود در بخش پنجم، آزمایش‌ها، مقایسه و نتایج موردبحث قرار گرفته است. در نهایت بخش ششم، نتیجه‌گیری و کارهای پیشرو و زمینه‌های تحقیقات جدید بررسی می‌شود.

۲- مروری بر مطالعات پیشین

شبکه‌های اجتماعی در حال حاضر یکی از جدیدترین موضوعات تحقیقاتی است. در این زمینه، به طور مثال مدل‌های رفتاری کاربران و یا ارزیابی‌های زمانی و منطقی ریاضی [۶]؛ و گره‌های حسگر در شبکه به‌منظور شناسایی گره‌های مخرب مطرح شده است [۷]. توجه به سوابق کاربران و یا مدل‌های مبتنی بر داده‌کاوی و شناسایی الگو نیز ارائه شده است [۵]. مدل‌های بررسی ویژگی‌های کاربران از جمله صمیمیت و شباهت و شهرت نیز از موضوعات مورد تحقیق بوده است [۸]. از روش‌های دیگر، تحلیل و ارزیابی گراف کوچکی از کاربران به‌منظور یافتن الگو و مدل کلی برای همه کاربران بوده است [۹]. تشکیل تیم‌هایی بر اساس روابط اعتماد میان کاربران در شبکه‌های اجتماعی برخط، که در آن اعتماد بر اساس تخصص و سطح اولویت حفظ حریم خصوصی کاربر اندازه‌گیری می‌شود [۱۰]. ولی عمده تحقیقات ارائه مدلی کلی است. و عمدتاً مدل‌های ارائه‌شده بر پایه محوریت کلی کاربران است و کمتر به منطق ذهنی و الگوی شناسایی فردی اعتماد پرداخته شده‌است. در این بخش به بررسی مدل‌های ارائه‌شده قبلی که با مبحث تحقیقی این مقاله تطابق دارد می‌پردازیم. پژوهش‌های مورد بررسی به دو دسته زیر تقسیم می‌شوند:

۲-۱- مطالعات با استفاده از نظریه دمپستر شفر

در سال ۲۰۱۵ لارنس چلوی، به ارائه راهکاری به‌منظور ارزیابی اعتماد و میزان مقبولیت آن توسط روابط منطقی پرداخته است. در این روش از روابط منطقی و ترکیب آن با نظریه دمپستر شفر برای تعیین قطعیت و تعیین صحت شواهد برای ارزیابی مدیریت اعتماد استفاده شده است. مسئله اصلی این است که چگونه یک کاربر به‌کاربر دیگر اعتماد و اطمینان پیدا کند و چگونه اطلاعات دریافتی از آن را باور کند. در این مدل ارزیابی می‌تواند هم از منبع مستقیم و هم از منبع غیرمستقیم باشد و در هر دو حالت با واسطه و بی‌واسطه این روش کارایی دارد که این از مزایای اصلی ارزیابی اعتماد در این مدل است. این الگوریتم شش حالت مختلف را برای تعیین اصالت اطلاعات در نظر می‌گیرد و یک چهارچوب منطق ریاضی برای ارزیابی اعتماد منابع استفاده می‌کند. این الگوریتم می‌تواند برای نتیجه‌گیری‌های اطلاعات پیچیده‌تر، به‌ویژه برای مسئله انتشار اعتماد مورد استفاده قرار گیرد و این زمانی رخ می‌دهد که کاربران، ارزیابی اعتماد خود را با یکدیگر مبادله و گزارش می‌کنند [۶].

ابزاری برای تصمیم‌گیری موردتوجه قرار گرفته است که در سال‌های اخیر توجه بیشتری را به خود جلب کرده است [۳]؛ و به‌طور گسترده در زمینه‌های مختلف مانند سیستم‌های توصیه‌گر، مشکلات تصمیم‌گیری، سیستم‌های همتا به همتا، اینترنت اشیا و شبکه‌های اجتماعی استفاده می‌شود [۱۱]. تعیین خصوصیات کاربران و ارزیابی آن‌ها می‌تواند راهکار مناسبی برای تعیین صداقت کاربران و میزان اعتماد به آنها می‌باشد. از جمله اهداف مهم، پیدا کردن روابط بین ویژگی‌های کاربران و ارزیابی اعتماد و تصمیم آن است. در اعتماد مستقیم در نظر گرفتن خصوصیات فردی کاربران و مدل‌های ذهنی، راهکار بهتر و مطمئن‌تری برای ارزیابی اعتماد است. در صورتی‌که در اعتماد غیرمستقیم زمانی شکل می‌گیرد که یک کاربر نمی‌تواند اعتماد مستقیم برای کاربر دیگر، را تحلیل و محاسبه کند که در این حالت از ارزیابی کاربران دیگر کمک می‌گیرد و در نهایت اجتماع اعتماد بر اساس روش میانگین وزن‌ها و در توافق با اعتماد مستقیم و غیرمستقیم به‌دست می‌آید [۲].

در مجموع، مدیریت اعتماد شامل سیستم‌ها و روش‌های محاسبه و تصمیم‌گیری افراد درباره میزان قابل‌اتکا بودن تراکنش‌های دربردارنده ریسک است و از جنبه دیگر به معنی فعالیت‌های جمع‌آوری، کدگذاری، تحلیل و ارزیابی مدارک مرتبط با شایستگی، صداقت، امنیت و قابلیت اتکا باهدف انجام قضاوت‌ها و تصمیم‌گیری‌های مربوط به روابط اعتماد است [۴].

امروزه باتوجه به ارتباطات برخط و وجود بستر اینترنت و کاربرد آن در همه مراکز مالی و شرکت‌های تجاری، بحث امنیت این گونه مراکز حساس و مهم است. ارتباط امن و محافظت از اطلاعات از چالش‌های اصلی این مراکز است، به‌طوری‌که بودجه زیادی توسط این شرکت‌ها صرف ایجاد بستر امن برای ارتباطات می‌شود و این در حالی است که بعضاً نقاط ضعف بسیاری وجود دارد و اعتماد اشتباه به کاربران متخلف، باعث ضرر و زیان زیادی به این مراکز شده است. در این بین ایجاد بستر امن و بالابردن ضریب اعتماد، کارساز است. این پژوهش، می‌تواند در شبکه‌های اجتماعی و نرم‌افزارهای حرفه‌ای ارتباطی، کاربرد مناسب داشته باشد و به شکل عملی باعث ارتقا ضریب امنیتی و ارتباطی این نرم‌افزارها و وبسایت‌ها باشد.

در زمینه اعتماد، آنچه در تحقیقات قبلی بیشتر مشهود است، استفاده از معیارهای ارزیابی کلی برای همه کاربران است. در تحقیقات پیشین یک الگوی سراسری برای همه در نظر گرفته شده است؛ در حالی که مدیریت اعتماد در حوزه امنیت نرم تعریف می‌شود و شامل معیارهای پویا است و در بازه زمانی ارزیابی، معیارها متغیر هستند؛ لذا در نظر گرفتن معیارهای ثابت، دقت ارزیابی را کاهش می‌دهد [۵]؛ بنابراین بر اساس این مقاله، روشی که هر کاربر بتواند ارزیابی خاص خود را داشته باشد و تصمیم و ارزیابی اعتماد را به‌صورت فردی در نظر بگیرد با دقت بیشتری همراه است و باتوجه به اینکه معیارهای اعتماد اصولاً کیفی هستند و تصمیم بر روی معیارهای کیفی فاقد دقت لازم است، استفاده از نظریه شواهد و تبدیل معیارهای کیفی به کمی راهکار بهتری است. همچنین پردازش مقادیر و ارزیابی‌ها با استفاده از معادلات ریاضی چالش‌برانگیز است به‌طوری‌که می‌توان ارزیابی را با صحت بیشتری انجام داد و ویژگی‌های کاربران و رفتار آنان را به معیارهای قابل‌اندازه‌گیری تبدیل کرد. از طرفی رفتار کاربران و تعامل رفتاری و معیارهای کیفی و بعضاً زمانی در محاسبات میزان اعتماد مهم است. همچنین استفاده از نظریه مجموعه خشن فازی در محاسبه معیارهای وزنی وابسته به ویژگی‌های کاربران و پردازش و ارزیابی اعتماد بر اساس آنها، ضریب دقت بیشتری دارد. استفاده از ریاضیات نظریه شواهد و به‌دست‌آوردن مقادیر عددی و ترکیب آنها به‌منظور ارزیابی سطح اعتماد از ضرورت‌ها و تکنیک‌های مناسب در مدیریت اعتماد است.

از موارد بااهمیت دیگر در ارزیابی اعتماد، توجه به وجود ابهام و عدم قطعیت است که ممکن است رخ دهد. اغلب روش‌های حال حاضر ابهام و عدم قطعیت

الگوریتم‌های مبتنی بر گراف را حل کرده و در نتیجه پوشش و دقت اعتماد در شبکه‌های اجتماعی برخط را افزایش می‌دهد. در این مدل ابتدا همسایگان هر کاربر (با تعریف محله بر اساس فاصله اجتماعی) مشخص می‌شود. سپس، به‌منظور یافتن مسیرهای قابل‌اعتماد بین هر جفت کاربر در شبکه اعتماد، از یادگیری خودکار به‌جای BFS استفاده می‌شود. به‌منظور یادگیری همسایگان قابل‌اعتماد بعدی، به هر گره از شبکه اعتماد یک یادگیری خودکار اختصاص داده می‌شود. در نهایت، یک نمودار قابل‌اعتماد شامل کاربران و همسایگان آنها برای ارزیابی اعتماد ایجاد می‌شود [۱].

در سال ۲۰۲۲ عبدالغنی و همکاران، یک مدل اعتماد چندسطحی پویا و مقیاس‌پذیر جدید را به نام DSL-STM پیشنهاد دادند. DSL-STM به طور خاص برای محیط‌های IoT طراحی شده و برای توصیف و رفتار موجودیت‌های IoT از معیارهای چندبعدی استفاده می‌کند. در این روش امکان طبقه‌بندی کاربران، شناسایی انواع حملات و مقابله با آنها توسط یک روش مبتنی بر یادگیری ماشین انجام می‌گیرد. DSL-STM در نهایت از یک روش انتشار ترکیبی برای گسترش ارزش‌های اعتماد در شبکه، کاهش مصرف منابع و حفظ مقیاس‌پذیری و پویایی استفاده می‌کند [۱۵].

در سال ۲۰۲۲ امیری زرنندی و همکاران، با استفاده از اطلاعات اجتماعی از شبکه، یک سیستم مدیریت اعتماد مبتنی بر بلاک‌چین (LBTM) سبک‌وزن را برای اینترنت اشیا اجتماعی پیشنهاد دادند. با استفاده از فناوری بلاک‌چین در رویکرد پیشنهادی، یک راه‌حل غیرمتمرکز ارائه می‌کند که ارزیابی اعتماد را امکان‌پذیر می‌کند. چارچوب پیشنهادی در برابر دشمنان انعطاف‌پذیر است و قادر است گره‌های مخرب را پس از تحلیل قابلیت اعتماد شناسایی کند [۱۶]. در سال ۲۰۱۸ شونان، بر پایه اعتماد به محیط شبکه به‌عنوان یک سیستم خود ساماندهی و با استفاده از اصل خود ساماندهی انتشار اعتماد اجتماعی، یک روش ارزیابی اعتماد ژنتیکی با بهینه‌سازی خود و ویژگی‌های خانوادگی ارائه کرد. در این روش، عوامل ارزیابی اعتماد شامل زمان، آدرس IP، بازخورد رفتاری و اعتماد بصری است. ساختار داده‌های جدول ضبط دسترسی و جدول ضبط اعتماد به‌منظور ذخیره ارتباط بین گره‌های اجداد و گره‌های فرزندان طراحی شده است. الگوریتم جستجوی اعتماد ژنتیکی با شبیه‌سازی فرایند تکامل بیولوژیکی طراحی شده است. با جستجوی ژنتیکی، کروموزوم اعتماد بهینه در جمعیت انتخاب شده و ارزش اعتماد کروموزوم محاسبه می‌شود که نتیجه، ارزیابی اعتماد ژنتیکی گره است [۱۷].

در سال ۲۰۲۰ احمدیان و همکاران، روشی مبتنی بر شهرت کاربران با افزایش دقت رتبه‌بندی کاربرانی که رتبه‌بندی قابل‌قبول و اطلاعات اعتماد کافی ندارند، پیشنهاد دادند. در این روش به طور خاص، از قابلیت اطمینان کاربر برای تعیین اثربخشی رتبه‌بندی و شبکه‌های اعتماد کاربران استفاده شده است. برای این منظور، از مدل شهرت کاربر پیشنهادی برای پیش‌بینی رتبه‌بندی مجازی استفاده شده است. علاوه بر این، از تنوع، تازگی و قابلیت اطمینان در فرایند افزایش رتبه‌بندی پیشنهادی کاربران استفاده شده است [۱۸].

در سال ۲۰۲۰ ون مو و همکاران، یک رویکرد ارزیابی اعتماد فعال و قابل‌تأیید برای شناسایی اعتبار دستگاه‌های متصل به اینترنت ارائه دادند. در مقایسه با رویکرد اعتماد غیرفعال این مدل یک مکانیسم اعتماد فعال است و دارای مزیت در کسب اعتماد است [۱۹].

در سال ۲۰۲۰ نصیر و کیم، در رابطه با ارزیابی اعتماد در شبکه‌های اجتماعی الگوریتمی را ارائه دادند. در این روش ارزیابی مبتنی بر استناد مشترک اعتماد کاربران و انتشار اعتماد است. به طور متوسط تفاوت اعتماد دو کاربر نسبت به ارزیابی سایر کاربران مشخص می‌شود. با استفاده از این تفاوت، چهار میزان اعتماد جزئی تخمین زده می‌شود و ارزش اعتماد نهایی به‌عنوان میانگین وزنی این تخمین‌های جزئی محاسبه می‌شود [۲۰].

در سال ۲۰۱۸ ژانگ و همکاران، از ارزیابی مدیریت اعتماد بر مبنای نظریه شواهد دمپستر شفر یا همان تابع باور برای شناسایی گره‌های مخرب در شبکه‌های حسگر بی‌سیم استفاده کرده است. در این روش باتوجه به همبستگی فضایی بین داده‌های جمع‌آوری شده توسط گره‌های حسگر در ناحیه مجاور، می‌توان درجه اعتماد را محاسبه کرد. همچنین به‌منظور ارزیابی اطمینان و عدم‌اطمینان اعتماد، یک طرح مدیریت اعتماد جدید بر مبنای نظریه شواهد دمپستر شفر برای تشخیص گره‌های مخرب پیشنهاد شده است. در این مدل از طریق شمارش تعداد رفتارهای تعاملی بین گره‌ها، برای ارزیابی اعتماد مستقیم و اعتماد غیرمستقیم استفاده شده است [۷].

در سال ۲۰۲۰ بهشید شایسته و همکاران، روشی برای محاسبه اعتماد ارائه دادند که علاوه بر محاسبه اعتماد موجودیت‌ها در یک کاربرد اینترنت اشیا، به محاسبه اعتماد‌پذیری داده نیز می‌پردازد. روش پیشنهادی آن‌ها از یادگیری بیزی برای محاسبه اعتماد با دنظرگرفتن ارتباط بین اعتماد‌پذیری داده و اعتماد موجودیت استفاده می‌کند و برای محاسبه اعتماد‌پذیری داده از قانون ترکیب نظریه دمپستر - شفر استفاده می‌کند [۱۱].

در سال ۲۰۱۰ زیانگ و ژانگ، الگوریتمی در زمینه مدیریت اعتماد و تابع باور و انتقال اعتماد بر پایه نظریه دمپستر شفر ارائه کردند. در این مدل دو نوع رابطه اعتماد نمایش داده می‌شود که یکی شامل هویت اعتماد و دیگری رفتار اعتماد است که این دو در اعتماد مستقیم و اعتماد غیرمستقیم تعریف می‌شوند، سپس یک شبکه انتقال اعتماد و انتشار اعتماد و ترکیب قوانین ایجاد می‌شود. در نهایت، یک روش برای تبدیل سه‌گانه نظریه شواهد به یک نتیجه ساده برای انتخاب موجودیت اعتماد استفاده می‌شود و با استفاده از درجه نزدیکی، برای تجزیه و تحلیل و انتقال اعتماد استفاده می‌شود. معادلات به‌دست‌آوردن مقادیر ترکیب شواهد و انتقال و تجمیع اعتماد در این مدل مشاهده می‌شود [۱۳].

در سال ۲۰۱۸ جانانی و همکاران، روشی به‌منظور محاسبه اعتماد توزیع شده و تشخیص رفتار غیرعادی کاربران با استفاده از قضیه بیزین و نظریه شواهد دمپستر شفر ارائه کردند [۱۴].

۲-۲- مطالعات بدون استفاده از نظریه دمپستر شفر

در سال ۲۰۲۳ یان وی شو و همکاران، یک مدل شبکه‌های عصبی مبتنی بر توجه جدید برای پیش‌بینی اعتماد، به نام GainTrust، پیشنهاد داد. GainTrust با طراحی لایه جاسازی شبکه نا همگن، همسایگان اعتماد کاربران و سوابق رفتاری آن‌ها را در یک فضای ویژگی مشترک ترسیم می‌کند. در ادامه، از یک شبکه عصبی چندلایه LSTM برای یادگیری ویژگی‌های سری زمانی عملکرد کاربران در طول چندین بازه زمانی استفاده می‌کند. در انتها، پس از ادغام ویژگی‌های سری زمانی با ویژگی‌های محله مورداعتماد، با طراحی یک مکانیزم دوسطحی چند سر Attention، ویژگی‌های اعتماد جهانی برای ارزیابی اعتماد بین کاربران به دست می‌آید [۲].

در سال ۲۰۲۰ مریم فیاض و همکاران، از محاسبات نرم و مدل شبکه عصبی، برای پیش‌بینی مقادیر اعتماد استفاده کرده‌اند. این روش قابل‌اجرا برای شبکه‌های اعتماد مختلف با مقادیر متفاوت اعتماد است. برای استنتاج اعتماد، چهار ویژگی از شبکه اجتماعی استخراج می‌شود که در برگزیده تأثیرات جوانب متفاوت اعتماد بر فرایند تصمیم‌گیری است. به‌علاوه، از الگوریتم ژنتیک برای تنظیم وزن‌ها و بایاس شبکه عصبی استفاده شده است. روش پیشنهادی به‌عنوان یک ابزار قدرتمند در استخراج روابط و استنتاج اعتماد در یک شبکه اجتماعی مورد استفاده قرار گرفته است [۱۲].

در سال ۲۰۲۲ فاتحی و همکاران، با ترکیب رویکرد مبتنی بر گراف و رویکرد هوش مصنوعی، مدلی را پیشنهاد دادند که با یافتن تمام مسیرهای قابل‌اطمینان بین کاربران به طور قابل‌توجهی چالش طول مسیرها در

باشد که میزان اهمیت هر کدام از شهود را بیان می‌کند. معادلات اصلی نظریه شواهد به صورت زیر است:

$$m(\emptyset) = 0 \quad (۱)$$

$$\sum_{A \subseteq \Omega} m(A) = 1 \quad (۲)$$

$$Bel: 2^{\Omega} \rightarrow [0,1] \quad (۳)$$

$$\begin{aligned} Bel(\emptyset) &= 0 \\ Bel(\Omega) &= 1 \end{aligned} \quad (۴)$$

$$Bel(A) = \sum_{X \subseteq A} m(X) \quad (۵)$$

۳-۲- مجموعه خشن فازی

نظریه مجموعه‌های ناهموار فازی [۲۵] بر دو نظریه دیگر بنا شده است، یعنی نظریه مجموعه‌های فازی [۲۶] و نظریه مجموعه‌های خشن [۲۷]. مجموعه خشن سعی می‌کند جهان گفتمان را به ناحیه مثبت (مجموعه‌های تقریب پایین‌تر)، ناحیه مرزی و ناحیه منفی تقسیم کند. مجموعه خشن فازی برای هر نمونه یک جفت عضویت را برمی‌گرداند که عضویت تقریبی پایین‌تر را به عنوان درجه قطعیت و عضویت تقریبی بالایی را به عنوان درجه احتمالی برای گنجاندن در نمونه‌های هدف نشان می‌دهد.

عضویت‌های تقریبی پایین و بالایی توسط روابط غیرقابل تشخیص (IR) بین نمونه‌ها و میزان وابستگی به مجموعه هدف، μ_F ساخته می‌شوند [۲۸]. IR همانطور که در (۶) نشان داده شده است، شباهت هر جفت نمونه را اندازه می‌گیرد. هنگامی که دو نمونه یکسان هستند، IR برابر با یک می‌شود و زمانی که نمونه‌ها کاملاً متفاوت هستند، صفر را نشان می‌دهد.

$$IR(x_i, x_j) = \tau_{\alpha \in Dimension} (1 - |x_i(\alpha) - x_j(\alpha)|^2) \quad (۶)$$

که در آن τ یک هنجار مثلثی است (t-norm)، $[0, 1] \rightarrow [0, 1]^2$ بر اساس بعد (ویژگی) α .

باتوجه به رابطه فوق، تفاوت بین دو نمونه در هر بعد با یک t-norm جمع می‌شود. نتیجه رابطه غیرقابل تشخیص بین دو نمونه تحت مرز تصمیم DB نامیده می‌شود. از سوی دیگر، عضویت هر نمونه در کلاس هدف، μ_F می‌تواند به طور متفاوتی مانند نمونه‌های باینری یا سایرین نشان داده شود. بنابراین، با این دو مفهوم، عضویت‌های تقریبی پایین‌تر و بالا به ترتیب در (۷) و (۸) تعریف می‌شوند:

$$\mu_{F_{IR, \mu_F}}(x_i) = \inf_{x_j \in T, x_i \neq x_j} I(IR(x_i, x_j), \mu_F(x_j)) \quad (۷)$$

$$\mu_{\overline{F_{IR, \mu_F}}}(x_i) = \sup_{x_j \in T, x_i \neq x_j} \tau(IR(x_i, x_j), \mu_F(x_j)) \quad (۸)$$

همان‌طور که در معادلات بالا نشان داده شده است، مقادیر $\mu_{F_{IR, \mu_F}}$ و $\mu_{\overline{F_{IR, \mu_F}}}$ عضویت‌های تقریبی پایین و بالا را نشان می‌دهد. عملگرهای فازی، $implicator(I)$ و $t-norm(\tau)$ ، دو عنصر اساسی را با هم ترکیب کرده و چندین خروجی را به وجود می‌آورند. سپس «inf» و «sup» یکی از این نتایج را به عنوان نتیجه نهایی انتخاب می‌کند. از این رو، داده‌های پرت و نویز می‌توانند عضویت تقریب پایین‌تر و بالایی را در محدوده وسیعی تغییر دهند. بنابراین، Verbiest و همکاران [۲۸] نسخه‌های تعدیل‌شده عضویت‌ها را با استفاده از میانگین وزنی سفارشی (OWA) به جای «inf» و «sup» پیشنهاد می‌کند. این اشکال عضویت به صورت زیر ارائه می‌شوند:

$$\mu_{F_{IR, \mu_F}}(x_i) = OWA_{min} I(IR(x_i, x_j), \mu_F(x_j)) \quad (۹)$$

در سال ۲۰۲۱ وانگ و تونگ، روشی را ارائه دادند که میزان اعتماد کاربران به یکدیگر در یک شبکه اجتماعی با یک روش جدید یادگیری چهارگانه نامتقارن (به جای روش سه‌گانه نامتقارن کلاسیک) محاسبه می‌شود. این روش با استفاده از داده‌های نمونه ارزیابی شده و سپس با روش انتقال یادگیری، ارزیابی اعتماد در یک شبکه اجتماعی جدید پیش‌بینی می‌شود [۲۱].

در سال ۲۰۲۰ سعیدی، چهار روش تخمین میزان اعتماد و تحلیل آن توسط فرضیه‌های آماری کولموگورف - اسمیرنوف و اندرسون - دارلینگ برای شناسایی و اعتبارسنجی بهترین مدل، پیشنهاد داد. در این روش از یک الگوریتم ابتکاری مبتنی بر روش کلونی زنبورها، به منظور بهینه‌سازی استفاده شده است [۲۲].

در سال ۲۰۲۱ وو و تیان، بر اساس ترکیب شبکه‌های پتری فازی و تحلیل رفتاری کاربران به منظور ارزیابی اعتماد، الگوریتمی را پیشنهاد دادند. در این الگوریتم ابتدا از شبکه‌های پتری فازی برای محاسبه میزان اعتماد پیشنهادی کاربر به عنوان میزان اعتماد غیرمستقیم و از سوابق رفتاری کاربران به عنوان شهودی برای ارزیابی اعتماد مستقیم استفاده شده است. در نهایت، ترکیب دو ارزیابی مستقیم و غیرمستقیم میزان اعتماد جامع کاربران را مشخص می‌کند [۲۳].

۳- مفاهیم پایه

در این بخش به بیان و تعریف مفاهیم پایه‌ای در رابطه با ارزیابی اعتماد و نظریه شواهد دمپستر شفر و مجموعه خشن فازی می‌پردازیم.

۳-۱- نظریه شواهد دمپستر شفر

نظریه شواهد دمپستر شفر یک ابزار بسیار مفید برای تصمیم‌گیری در شرایط عدم قطعیت و وجود ابهام است. تابع باور یا همان نظریه شواهد نگارشی از شهود غیرقابل اندازه‌گیری در مقیاس قابل‌سنجش ارائه می‌دهد و با محاسبه و اندازه‌گیری و رفع ابهام به مقادیر مشخص برای تصمیم‌گیری قطعی می‌رسد که این مهم‌ترین بخش کاربردی نظریه شواهد است. در این نظریه کلیه تصمیم‌ها در یک مجموعه قرار می‌گیرند که هر کدام از این حالت‌ها یک عضو کانونی و اصلی است و البته می‌توان ترکیبی از این حالت‌ها را در نظر گرفت. به عنوان مثال اگر در رابطه با عملی سه عامل مختلف را داشته باشیم می‌توان برای هر کدام یک احتمال عمل انجام شده را در نظر گرفت:

جدول ۱- شواهد و باور

$m(A) = 0.1$	$m(A) = 0.15$
$m(B) = 0.3$	$m(B) = 0.05$
$m(C) = 0.05$	$m(C) = 0.05$
$m(A, B) = 0.2$	$m(A, B) = 0.1$
$m(A, C) = 0.3$	$m(A, C) = 0.1$
$m(B, C) = 0.1$	$m(B, C) = 0.1$
$m(\theta) = 0.05$	$m(\theta) = 0.45$

مثال بالا شامل حالت‌های مختلفی است که از این اعضای کانونی می‌توان در نظر گرفت که احتمالات مختلف برای اعضای کانونی و ترکیبی متصور است. ترکیب شواهد می‌تواند بر اساس قوانین خاص در نظر گرفته شده، انجام می‌گیرد. البته در ترکیب شواهد در برخی موارد ممکن است تداخل رخ دهد که در این مواقع استفاده از محاسبه شباهت و اندازه‌گیری ابهام برای به دست آوردن وزن و اعتبار شواهد کارساز است. بدین شکل می‌توان از تداخل شواهد و اشتباه در تصمیم‌گیری جلوگیری کرد. استفاده از این وزن‌ها در تعیین مقدار اهمیت شواهد و درستی آنها تصمیم نهایی را دقیق‌تر می‌کند. در نهایت ممکن است همراه با مجموعه شواهد، وزن آن‌ها نیز به صورت یک مجموعه کامل وجود داشته

فردیت شکل می‌گیرد، اعتماد فردی نامیده می‌شود. میزان اعتماد فردی شامل ارزیابی دو مجموعه است که یکی مجموعه وزن‌ها و دیگری مجموعه محدوده‌های مقادیر ویژگی‌هاست. مجموعه وزن‌ها اهمیت ویژگی‌ها در ارزیابی اعتماد است. مجموعه ویژگی‌ها شامل درجه فعالیت و دنبال‌کننده و گروه و شهرت است. مجموعه مقادیر محدوده ویژگی‌ها شامل بازه کمینه و بیشینه مقادیر ویژگی‌ها است. این محدوده‌ها از کاربران قابل‌اعتماد و غیرقابل‌اعتماد نمونه‌های داده شده، استخراج می‌شوند. اگر این محدوده‌ها جدا از هم باشند طبیعتاً تداخل وجود ندارد؛ ولی اگر این محدوده‌ها در قسمتی، اشتراک داشته باشند، در این حالت بازه اشتراکی به‌عنوان محدوده ابهام تعریف می‌شود. به‌ازای محدوده مقادیر هر ویژگی و میزان شباهت آن با محدوده‌های ویژگی، بخشی از شواهد اعتماد تولید می‌شوند و در کل تمام ویژگی‌ها با مقادیر مشخص، شواهد اعتماد را تولید می‌کنند و ترکیب این شواهد با توجه به وزن و اهمیت آنها بیانگر آمادگی نهایی برای تصمیم اعتماد است. با توجه به ترکیب شواهد و به‌منظور تصمیم نهایی می‌توان حالات زیر را در نظر گرفت:

$$m(Trust) < m(Distrust) \quad (15)$$

$$m(Trust) \geq m(Distrust) \quad (16)$$

در این مدل با محاسبه میانگین مقادیر اعتماد کاربران و معیار اطمینان، میزان اعتماد هر کاربر در کل شبکه اجتماعی محاسبه می‌شود و به‌عنوان یک اعتبار و معیار کلی در نظر گرفته می‌شود.

۴-۱- الگوریتم محاسبه مقادیر ویژگی‌ها

درجه فعالیت بر اساس تعداد نظرات و پست‌ها در یک بازه زمانی مشخص نسبت به کل پست‌ها و نظرات محاسبه می‌شود. درجه فعالیت با استفاده از رابطه زیر محاسبه می‌شود.

$$Activity Degree = \frac{(N_{comment}^n + N_{post}^n)}{n} \quad (17)$$

جایی که n تعداد کل پست‌ها و نظرات، $N_{comment}^n$ تعداد نظرات و N_{post}^n تعداد پست‌ها در یک بازه زمانی مشخص می‌باشد.

درجه تعداد دنبال‌کننده به نسبت اشتراک به اجتماع تعداد دنبال‌کننده‌های دو کاربر گفته می‌شود که با استفاده از رابطه زیر محاسبه می‌شود:

$$Follower Degree = \frac{|Fl_i \cap Fl_j|}{|Fl_i \cup Fl_j|} \quad (18)$$

که در آن Fl_i تعداد دنبال‌کننده‌های کاربر i ، Fl_j تعداد دنبال‌کننده‌های کاربر j و علامت $||$ قدرمطلق می‌باشد.

درجه گروه و تگ‌ها به حاصل ضرب نسبت‌های اشتراک به اجتماع گروه‌ها عضو شده دو کاربر و تگ‌های آن‌ها گفته می‌شود و از رابطه زیر به دست می‌آید:

$$Group Degree = \frac{|Tg_i \cap Tg_j|}{|Tg_i \cup Tg_j|} \times \frac{|Gr_i \cap Gr_j|}{|Gr_i \cup Gr_j|} \quad (19)$$

که در آن Tg_i و Tg_j بترتیب تگ‌های کاربر i و کاربر j و Gr_i ، Gr_j بترتیب گروه‌های عضو شده کاربر i و کاربر j می‌باشد.

درجه شهرت بر اساس میانگین درجه کاربران دیگر به کاربر مورد ارزیابی و کاربر ارزیابی‌کننده محاسبه می‌شود و از رابطه زیر به دست می‌آید:

$$Reputation Deg = \frac{(Avg(R_i) + Avg(R_j))}{2} \quad (20)$$

در رابطه ۲۰ درجه کاربران با R نشان داده شده است.

۴-۲- الگوریتم محاسبه مقادیر اعتماد بر اساس مقادیر ویژگی‌ها

در ابتدا مجموعه داده‌های آزمون دو مجموعه کاربران معتمد u_t و غیر معتمد

$$\mu_{\overline{FIR}, \mu_F}(x_i) = OWA_{max} \tau(IR(x_i, x_j), \mu_F(x_j)) \quad (10)$$

لیو و همکاران [۲۹، ۳۰، ۳۱] در واقع سعی بر آن است تا احتمال وقوع زیرمجموعه‌ای از جهان گفتمان را بر اساس توابع باور و معقول بودن به ترتیب در قالب (۱۱) و (۱۲) توضیح دهد.

$$\forall A \in T: Bel(A) = \sum_{B: B \subseteq A} P_\alpha(B) \quad (11)$$

$$\forall A \in T: PI(A) = \sum_{B: B \cap A \neq \emptyset} P_\alpha(B) \quad (12)$$

جایی که $P_\alpha(B)$ احتمال وقوع B را نشان می‌دهد در حالی که مقدار α احتمال وقوع مربوط به $\alpha - cut$ در مفهوم فازی را دارد. از سوی دیگر، رابطه بین این دو مفهوم به شرح زیر است:

$$\forall A \in T: Bel(A) = \mu_{\overline{FIR}, \mu_F}(A) \quad (13)$$

$$\forall A \in T: PI(A) = \mu_{\overline{FIR}, \mu_F}(A) \quad (14)$$

در نتیجه، مجموعه خشن فازی برای مدیریت ابهام و تضاد بین داده‌ها قدرتمند است.

۳-۳- ویژگی‌های منتخب کاربران

در شبکه‌های اجتماعی رفتار و ویژگی‌های کاربران از موارد بااهمیت در تعیین صداقت و اعتماد کاربران است. در این زمینه می‌توان ویژگی‌های مختلفی از کاربران را در نظر گرفت؛ ولی باید توجه داشت بعضی از این ویژگی‌ها به شکل درست قابل‌ارزیابی و محاسبه نیستند و مواردی را می‌توان در نظر گرفت که از قابلیت محاسباتی بیشتری برخوردار باشند. مجموعه ویژگی‌هایی که در این مقاله در نظر گرفته شده است شامل درجه فعالیت و درجه تعداد دنبال‌کننده و درجه گروه و درجه شهرت است و این چهار ویژگی، ویژگی‌های اصلی و بااهمیت را تشکیل می‌دهند. درجه فعالیت بر اساس تعداد نظرات و پست‌ها نسبت به یک بازه انتخابی برای این ویژگی است. درجه تعداد دنبال‌کننده توسط اشتراک مجموعه‌های دنبال‌کننده دو کاربر تقسیم بر اجتماع دو مجموعه به دست می‌آید. درجه گروه و تگ‌ها بر اساس حاصل ضرب نسبت‌های اشتراک به اجتماع گروه‌ها عضو شده دو کاربر و تگ‌های آنان به دست می‌آید. درجه شهرت بر اساس نسبت امتیازاتی که بقیه کاربران داده‌اند و امتیازی که به‌وسیله‌ی کاربر تحلیلگر داده شده است، تعریف می‌شود.

در سه ویژگی اول با توجه به ماهیت آنها بازه‌ی مقدار ویژگی بین صفر و یک است یعنی درجه فعالیت و درجه دنبال‌کننده و درجه گروه حداکثر یک می‌باشد ولی درجه شهرت با توجه به وضعیت آن به صورت بازه انتخابی تعریف می‌شود که بیشینه آن بهترین درجه در شبکه اجتماعی مذکور می‌باشد.

۴- روش پیشنهادی

برای ارزیابی اعتماد در روش پیشنهادی از ویژگی‌های کاربران استفاده می‌شود. ارزیابی اعتماد بر اساس خصوصیات و واکنش‌ها و رفتار فردی کاربران است. در این مدل ابتدا ویژگی‌های قابل‌محاسبه کاربران اندازه‌گیری می‌شود و محدوده ویژگی‌ها مشخص می‌شود و با محاسبه شباهت مقادیر ویژگی‌ها، شواهد اعتماد تولید می‌شوند و با ترکیب وزنی این شواهد مقادیر اعتماد و بی‌اعتمادی و ابهام مشخص می‌شود و سپس با استفاده از این معیارها، اعتماد کلی کاربران محاسبه می‌شود. خروجی مدل تصمیم نهایی اعتماد خواهد بود.

باتوجه به بخش مطالعات پیشین، دیده شده که هر کاربر به ویژگی‌های مختلف اهمیت می‌دهد و برای تصمیم اعتماد ممکن است ویژگی‌های خاصی را در ارزیابی خود ملاک قرار دهد و به همین جهت این نوع اعتماد را که بر اساس

u_{dt} مشخص می‌شوند:

$$u_{dt} = \{u_{dt1}, u_{dt2}, \dots\} \quad (23)$$

در روابط فوق us مجموعه داده‌های آزمون u_t کاربران معتمد و u_{dt} کاربران غیر معتمد می‌باشند.

سیس در هر دو مجموعه کاربران معتمد و غیرمعتمد، بیشینه و کمینه مقادیر تمام ویژگی‌ها استخراج می‌شوند:

$$F: \{f_1, f_2, f_3, f_4\} \\ \forall f \in F \quad (24)$$

$$Sco_f = (\min(val(f), \max(val(f))) \quad (25)$$

در رابطه ۲۴ ویژگی‌ها با f_n نمایش داده شده اند و Sco_f مقدار بیشینه و کمینه ویژگی می‌باشد.

در ادامه محدوده ویژگی‌های اعتماد و بی‌اعتمادی توسط بیشینه و کمینه ویژگی‌ها مشخص می‌شود:

$$Sco_t = (\min(f_t), \max(f_t)) \quad (26)$$

$$Sco_{dt} = (\min(f_{dt}), \max(f_{dt})) \quad (27)$$

در رابطه ۲۶ محدوده ویژگی اعتماد با Sco_t و محدوده ویژگی بی‌اعتمادی با Sco_{dt} نشان داده شده است.

محدوده ابهام هر ویژگی نیز بر اساس اشتراک دو محدوده اعتماد و بی‌اعتمادی تشکیل می‌شود و با Sco_{am} نشان داده شده است:

$$Sco_{am} = Sco_t \cap Sco_{dt} \quad (28)$$

$$Sco_f = \langle Sco_t, Sco_{dt}, Sco_{am} \rangle \quad (29)$$

بر اساس مقادیر ویژگی‌های کاربران معتمد و غیرمعتمد و ارزیابی آنها، مقادیر وزن ویژگی‌ها به دست می‌آید. برای محاسبه وزن ویژگی‌ها از الگوریتم فازی رافست استفاده شده است.

برای این منظور، ts که فاصله معکوس نمونه‌ها از مرکز کلاس را نشان می‌دهد با استفاده از رابطه زیر محاسبه می‌شود:

$$t_s = 1 - \frac{\sum_{\alpha \in dimension} |\sum_{x_i \in T} x_i(\alpha) - c^*|^2}{|c^*|} \quad (30)$$

که در آن T مجموعه هدف، α بعد (ویژگی) و $\|$ مقدار مطلق است. همچنین C^* مرکز کلاس است. مرکز هر کلاس C^* به صورت زیر تعریف می‌شود:

$$c^* = \frac{1}{N} \sum_{x_i \in y_i} x_i \quad (31)$$

که در آن N تعداد کل نمونه‌های آموزشی در هر کلاس است. بنابراین، عضویت‌های تقریبی پایین و بالا مطابق (۷) و (۸) توسط موارد زیر بازسازی می‌شوند.

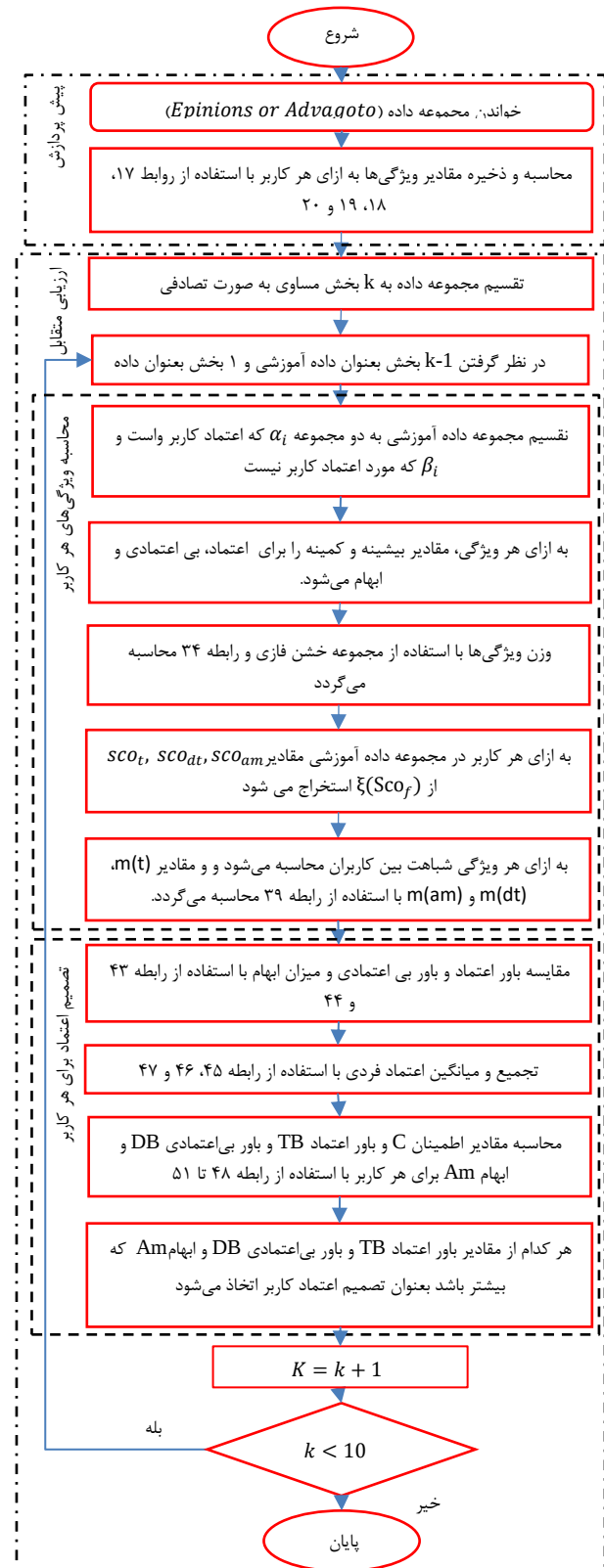
$$\mu_{FIR, \mu F}(x_i) = \inf I(IR(x_i, c^*), t_s) \quad (32)$$

$$\mu_{\overline{FIR}, \mu F}(x_i) = \sup \tau(IR(x_i, c^*), t_s) \quad (33)$$

عضویت‌های تقریبی در ساخت ED از (۶) و (۳۰) به‌عنوان IR و مقدار نمونه‌های متعلق به مجموعه هدف t_s استفاده می‌کنند.

$$ED(x_i) = - \sum_{x_i \in R^d} \mu_{\overline{FIR}, \mu F}(x_i) \log_2 \mu_{FIR, \mu F}(x_i) \quad (34)$$

بنابراین، در این مقاله، ED نتیجه می‌گیرد که هر دو عضویت تقریب پایین و بالا به‌عنوان شاخص قطعیت پیشنهاد شده‌اند. در واقع، ED یک چشم‌انداز کلی از نقش‌های نمونه با عضویت تقریبی بالا به دست می‌آورد. به دنبال این، عضویت تقریبی پایین‌تر که رابطه محدود بین نمونه‌ها و مرز تصمیم را محاسبه می‌کند، به ED اضافه می‌شود تا سطح نمونه‌ها را که مطمئناً در مرز تصمیم قرار



شکل ۱- نمودار جریان

$$us = \{(u_1, t_1), (u_2, t_2), \dots\} \quad (21)$$

$$u_t = \{u_{t1}, u_{t2}, \dots\} \quad (22)$$

اما میزان اعتماد هر کاربر به کاربر دیگر مشخص شده است و با تجمیع و میانگین می‌توان میزان نهایی اعتماد فردی قابل‌پذیرش برای هر کاربر را در شبکه مشخص کرد:

$$Trust = Avg\ m(T) \quad (45)$$

$$Distrust = Avg\ m(Dt) \quad (46)$$

$$Ambiguity = Avg\ m(Am) \quad (47)$$

باتوجه به خروجی مدل پیشنهادی مقادیر اطمینان، باور اعتماد، باور بی اعتمادی و ابهام Am تعریف می‌شوند که با رابطه‌های زیر معیار دقیق اعتماد و اطمینان آن را می‌توان مشخص کرد:

$$C = 1 - Am \quad (48)$$

$$TB = C \times (1 - T) \quad (49)$$

$$DB = C \times (1 - Dt) \quad (50)$$

$$TB + DB + Am = 1, \rightarrow Am = 1 - (TB + DB) \quad (51)$$

در روابط بالا مقادیر اطمینان با C ، باور اعتماد با TB ، باور بی‌اعتمادی با DB و ابهام با Am تعریف می‌شوند.

در این روش ارزیابی اعتماد فردی محاسبه می‌شود. آنچه مشخص است، ارزیابی اعتماد توسط مقادیر عددی ویژگی‌ها است که توسط الگوریتم مدل به دست می‌آید. تصمیم اعتماد باتوجه به پارامترهای اعتماد و بی‌اعتمادی و ابهام تعیین می‌شود. این مقادیر با کمک نظریه شواهد و رابطه‌های ریاضی به دست می‌آیند.

Algorithm: Trust decision based on user characteristics

Input: Epinions or Advagoto dataset

Output: Trust decision for users

```

1: for each user do
2:   Calculating the values of user attributes using the
     relationship 17, 18, 19 and 20
3:   Saving the attribute values of each user in a matrix
4:   Apply k – fold on the data set for k = 10
5:   for each k do
6:     Considering k – 1 fold as a training set and
       1 fold as a validation set
7:     Tell seti apart into  $\alpha_i$  and  $\beta_i$  where  $\alpha_i$  is user
       trusted by  $u_i$  and  $\beta_i$  is not trusted by  $u_i$ 
8:     for  $\forall f \in F$  do
9:       extract  $max_f$  and  $min_f$  of  $f$  from  $\alpha_i$ 
10:      extract  $max'_f$  and  $min'_f$  of  $f$  from  $\beta_i$ 
11:      set  $sco_t = (\min_f, \max_f)$ ,  $sco_{dt} = (\min'_f, \max'_f)$ 
12:      generate scope  $sco_{am}$  based on  $sco_t$  and  $sco_{dt}$ 
13:      set  $\xi_i(Sco_f) = \langle sco_t, sco_{dt}, sco_{am} \rangle$ 
14:      Generate evidence weights using relations
       34 and set  $g(\omega_1, \omega_2, \dots, \omega_f) = ED(u_i)$ 
15:      for each item  $\forall it \in set_i$  : do
16:        set  $ms = \emptyset$ 
17:        extract  $sco_t, sco_{dt}, sco_{am}$  from  $\xi_i(Sco_f)$ 
18:        for  $\forall f \in F$  do
19:          set the value of feature  $f$  belongs to
            it as  $val(f)$ 
20:          set the scope similarity between
             $[val(f), val(f)]$  and  $sco_t, sco_{dt}, sco_{am}$  as
             $sim_t, sim_{dt}, sim_{am}$ 
21:          set  $sum = sim_t + sim_{dt} + sim_{am}$ 
22:          set  $m(t) = (sim_t \times g) / sum$ ,
             $m(dt) = (sim_{dt} \times g) / sum$ ,
             $m(am) = (sim_{am} \times g) / sum$ 
23:   Comparison of trust belief and mistrust belief and
       amount of ambiguity:
        $m(T) < m(Dt)$  ,  $m(T) \geq m(Dt)$ 

```

دارند، بهبود بخشد. سپس، به محدوده $[0, 1]$ نگاشت می‌شود. از ED محاسبه شده به عنوان وزن ویژگی‌ها استفاده می‌شود.

$$W_i = ED(x_i) \quad (35)$$

در گام بعدی، مقدار شباهت بین مقدار هر ویژگی که به صورت بازه یکسان تعریف می‌شود، با محدوده‌های اعتماد و بی‌اعتمادی و ابهام به وسیله فرمول شباهت به دست می‌آید:

$$Sim_t, Sim_{dt}, Sim_{am} \quad (36)$$

$$Sim(a, b) = \frac{1}{1 + \varepsilon(a, b)}$$

$$\varepsilon(a, b) = abs\left\{\frac{1}{2}(a_{min} + a_{max}) + \frac{1}{2}(a_{max} - a_{min}) - \frac{1}{2}(b_{min} + b_{max}) - \frac{1}{2}(b_{max} - b_{min})\right\} \quad (37)$$

a و b دو محدوده از ویژگی‌ها هستند که رابطه فوق فاصله دو محدوده را محاسبه می‌کند. a_{min} و b_{min} به ترتیب کمترین مقادیر ویژگی دو محدوده و a_{max} و b_{max} نیز به ترتیب بیشترین مقادیر ویژگی دو محدوده مزبور است. برای سنجش میزان اختلاف دو محدوده، کوچکترین و بزرگترین اعداد دو محدوده در نظر گرفته شده است که در محاسبات میانگین کوچک‌ترین و بزرگ‌ترین اعداد محدوده و عرض دو محدوده براساس اختلاف بزرگ‌ترین و کوچک‌ترین اعداد محدوده در نظر گرفته شده است.

بعد از تعیین میزان شباهت، شهود نظریه شواهد دمپستر شفر بر اساس نسبت هر مقدار شباهت بر مجموع مقادیر شباهت برای هر ویژگی به دست می‌آید:

$$Sum = Sim_t + Sim_{dt} + Sim_{am} \quad (38)$$

$$m(t) = \frac{Sim_t}{Sum}, \quad m(dt) = \frac{Sim_{dt}}{Sum}, \quad m(am) = \frac{Sim_{am}}{Sum} \quad (39)$$

در روابط بالا Sim_t شباهت اعتماد، Sim_{dt} شباهت عدم اعتماد، Sim_{am} شباهت ابهام و Sum مجموع شباهت بین مقدار هر ویژگی می‌باشد. $m(t)$ مقدار شباهت اعتماد به مجموع مقادیر شباهت، $m(dt)$ مقدار شباهت عدم اعتماد به مجموع مقادیر شباهت، $m(am)$ مقدار شباهت ابهام به مجموع مقادیر شباهت می‌باشد.

در آخرین مرحله، شواهد اعتماد برای تمام ویژگی‌ها بر اساس وزن مربوطه با هم ترکیب می‌شوند و نهایتاً مقدار تابع جرم باور را نتیجه می‌دهد این عمل برای شواهد بی‌اعتمادی و ابهام نیز صورت می‌پذیرد و مقادیر تابع جرم بی‌اعتمادی و ابهام نیز محاسبه می‌شود:

$$m(T) = \frac{\sum w_i m_i(t)}{\sum w_i} \quad (40)$$

$$m(Dt) = \frac{\sum w_i m_i(dt)}{\sum w_i} \quad (41)$$

$$m(Am) = \frac{\sum w_i m_i(am)}{\sum w_i} \quad (42)$$

در روابط بالا، w_i وزن‌های ویژگی‌ها می‌باشد، $m(T)$ مقدار تابع جرم باور، $m(Dt)$ مقدار تابع جرم بی‌اعتمادی و $m(Am)$ مقدار تابع جرم ابهام می‌باشد. در این مرحله با مقایسه باور اعتماد و باور بی‌اعتمادی و مقدار ابهام و با استفاده از شاخص اطمینان ابهام می‌توان تصمیم اعتماد را گرفت و مقدار نهایی اعتماد فردی را مشخص کرد:

$$m(T) < m(Dt) \quad (43)$$

به این معنا که، کاربر u_i در مجموعه آموزشی به کاربر u اعتماد دارد.

$$m(T) \geq m(Dt) \quad (44)$$

به این معنا که، کاربر u در مجموعه آموزشی مورد اعتماد u_i نیست.

معیار F1-score، این معیار از میانگین هارمونیک دو معیار Recall و Precision تشکیل شده است و به صورت زیر فرموله می‌شود:

زمانی که می‌خواهید معیار ارزیابی شما میانگینی از دو معیار Recall و Precision باشد، می‌توانید از میانگین هارمونیک این دو معیار استفاده کنید که به آن معیار F1-score می‌گویند.

$$F1 - score = \frac{2}{\frac{1}{Recall} + \frac{1}{Precision}} = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (55)$$

۵-۳- الگوریتم پیشنهادی در مقایسه با سایر روش‌ها

روش پیشنهادی این پژوهش با روش‌های ارائه شده در مطالعات احمدیان و همکاران ۲۰۱۸، سعیدی ۲۰۲۰، عبدالغنی و همکاران ۲۰۲۲، فاتحی و همکاران ۲۰۲۲، مقایسه شده است [۱، ۱۵، ۲۲، ۱۸]. در مطالعه احمدیان و همکاران ۲۰۱۸، با استفاده از شهرت کاربران برای افزایش دقت رتبه‌بندی کاربرانی که رتبه‌بندی قابل قبول و اعتماد کافی ندارند روشی را پیشنهاد دادند. علاوه بر این، این روش از تنوع، تازگی و قابلیت اطمینان در فرایند افزایش دقت رتبه‌بندی پیشنهادی بهره می‌برد. در مطالعه سعیدی ۲۰۲۰، با استفاده از چهار روش تخمین میزان اعتماد و تحلیل آن توسط فرضیه‌های آماری کولموگورف - اسمیرنوف و اندرسون - دارلینگ برای شناسایی و اعتبارسنجی بهترین مدل، الگوریتمی پیشنهاد دادند در این روش از الگوریتم بهینه‌سازی کلونی زنبورها به منظور بهینه‌سازی اعتماد پیش‌بینی شده استفاده می‌شود. نتایج نشان‌دهنده برتری روش پیشنهادی این پژوهش در مقایسه با روش‌های ارائه شده در مطالعه احمدیان و همکاران و مطالعه سعیدی است. در مطالعه عبدالغنی و همکاران ۲۰۲۲، یک مدل اعتماد چندسطحی پویا و مقیاس پذیر جدید را پیشنهاد دادند که به طور خاص برای محیط‌های IIoT طراحی شده و از معیارهای چندبعدی استفاده می‌کند. در نهایت از یک روش انتشار ترکیبی برای گسترش ارزش‌های اعتماد در شبکه، کاهش مصرف منابع و حفظ مقیاس‌پذیری و پویایی استفاده می‌کند. در مطالعه فاتحی و همکاران ۲۰۲۲، با ترکیب رویکرد مبتنی بر گراف و رویکرد هوش مصنوعی، مدلی را پیشنهاد دادند که با یافتن تمام مسیرهای قابل اطمینان بین کاربران به طور قابل توجهی چالش طول مسیرها در الگوریتم‌های مبتنی بر گراف را حل کرده و در نتیجه پوشش و دقت اعتماد در شبکه‌های اجتماعی برخط را افزایش می‌دهد. در این روش به جای استفاده از BFS، به هر گره از شبکه اعتماد یک یادگیری خودکار اختصاص داده می‌شود. روش پیشنهادی ما در مقایسه با روش‌های ارائه شده در این دو مطالعه توانسته با دقت ۹۲،۵۴٪ به درستی در تصمیم‌گیری اعتماد عمل کند؛ نتایج نشان‌دهنده برتری روش پیشنهادی ما در مقایسه با این چهار مطالعه است.

در این پژوهش به منظور مقایسه و ارزیابی روش پیشنهادی با رویکردهای پیش گفته از روش «اعتبارسنجی متقابل»^۳ استفاده شده است. بدین منظور رویکرد پیشنهادی و چهار رویکرد [۱، ۱۵، ۲۲، ۱۸] در متلب پیاده‌سازی شده و با استفاده از اعتبارسنجی متقابل، در چهار معیار Recall، Accuracy، Precision و F1-score مقایسه و ارزیابی شده است. در اعتبارسنجی متقابل از روش k-Fold استفاده شده است. ابتدا مجموعه داده‌های آموزشی به طور تصادفی به k زیرنمونه یا «لایه»^۴ با حجم یکسان تفکیک شده است، سپس در هر مرحله از فرایند CV، تعداد k-1 از این لایه‌ها را به عنوان مجموعه داده آموزشی و یکی را به عنوان مجموعه داده اعتبارسنجی در نظر گرفته شده است. مقدار k برابر با ۱۰ در نظر گرفته شده است و فرایند CV، برای ۱۰ بار تکرار شده است. میانگین مقادیر بدست آمده از این ۱۰ تکرار برای هر معیار در هر

- 24: *Aggregation and average individual confidence:*
 $T = \text{Avg } m(T)$, $Dt = \text{Avg } m(Dt)$ and
 $Am = \text{Avg } m(Am)$
- 25: *Calculating the values:*
 $C = 1 - Am$, $TB = C \times (1 - T)$,
 $DB = C \times (1 - Dt)$, $Am = 1 - (TB + DB)$
- 26: $\forall f \in F$, set $\xi_i(f) = \text{Argmax } (TB, DB, Am)$
- 27: *Return ξ_i as the User - Will belongs to user u_i*

۵- آزمای‌ش‌ها

۵-۱- مجموعه داده‌ها

برای آزمای‌ش مدل ارائه شده از مجموعه داده اپینین استفاده می‌شود که یک مجموعه داده شناخته‌شده در حوزه شبکه‌های اجتماعی است. اغلب تحقیقات در حوزه اعتماد از این مجموعه داده استفاده می‌شود؛ زیرا دارای ساختار مناسب برای حوزه امنیت نرم است. از مجموعه داده ادوگاتو^۱ نیز به منظور ارزیابی و مقایسه مدل‌های اعتماد استفاده می‌شود. مجموعه داده اپینین^۲ دارای سه فایل اصلی است که در فایل اول داده‌های تصمیم اعتماد بین دو کاربر قرار دارد. در ستون My_ID کاربری که تصمیم اعتماد را گرفته است قرار دارد. ستون Other_ID شناسایی کاربری که اعتماد به او ارزیابی می‌شود قرار دارد و ستون Value ستونی است که مقدار تصمیم اعتماد قرار دارد که عدد یک برای اعتماد و عدد صفر برای بی‌اعتمادی در نظر گرفته شده است.

در فایل دوم مجموعه داده اپینین سه ستون وجود دارد که اولی Content_ID برای شناسایی نظر یا پست فرستاده شده به وسیله کاربر Author_ID است. که خود Author_ID نویسنده نظرات در ستون بعدی یعنی ستون دوم قرار دارد. در ستون بعدی Subject_ID قرار دارد که بیان‌کننده شناسایی موضوع نظری است که توسط نویسنده نظر فرستاده شده است. در فایل سوم هشت ستون قرار دارد که ستون Comment_ID بیان‌کننده شناسایی نظری است که توسط نویسنده فرستاده شده است و ستون دیگر در فایل سوم Object_ID است که شامل نظر یا پستی است که به وسیله کاربر با Member_ID امتیازدهی شده است که Member_ID نیز در ستون بعد قرار دارد و ستون دیگری که در فایل سوم اهمیت دارد ستون Rating است که از یک تا شش طبقه‌بندی شده است.

۵-۲- معیارها مورد ارزیابی

معیار Accuracy یا صحت، این معیار برابر است با تعداد مواردی که درست پیش‌بینی کردیم تقسیم بر تعداد کل پیش‌بینی‌هایی که انجام شده است و از رابطه زیر محاسبه می‌شود:

$$accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (52)$$

جایی که TP مخفف True Positive، TN مخفف True Negative، FP، False Positive و FN مخفف False Negative است.

معیار Recall یا یادآوری، این معیار به دنبال محاسبه‌ی پوشش بر روی کل داده‌هاست. به این معیار، معیار حساسیت (Sensitivity) نیز گفته می‌شود و به صورت زیر محاسبه می‌شود:

$$Recall = \frac{TP}{TP+FN} \quad (53)$$

معیار Precision یا دقت، در واقع این معیار به ما می‌گوید الگوریتم چند درصد «Positive» را درست پیش‌بینی کرده است. معیار دقت با استفاده از رابطه زیر محاسبه می‌شود:

$$Precision = \frac{TP}{TP+FP} \quad (54)$$

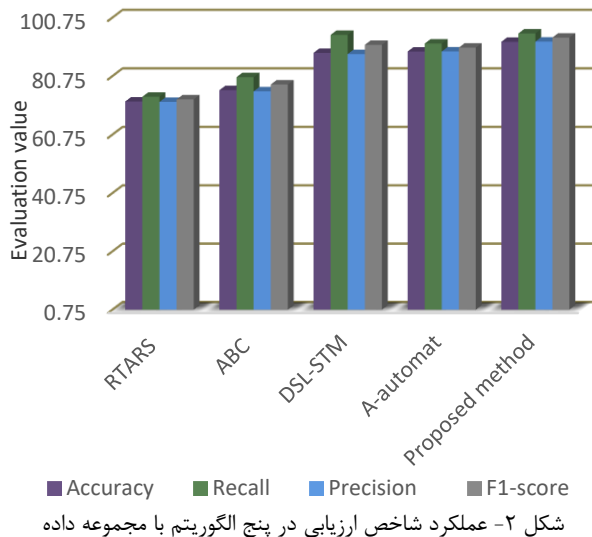
^۳ Cross Validation

^۴ Fold

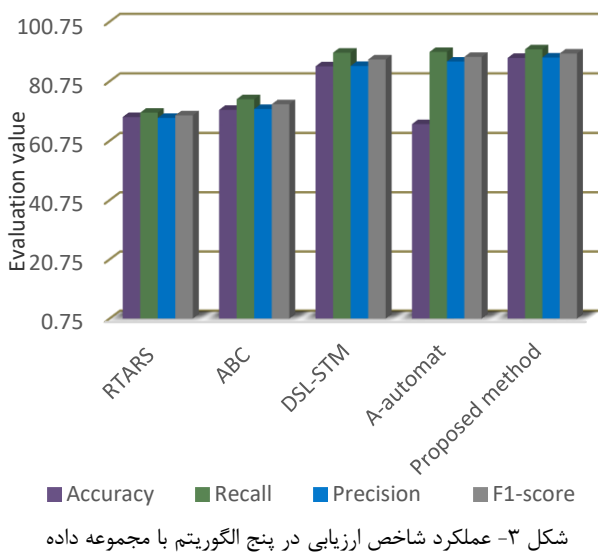
^۱ Advogato dataset

^۲ Epinions dataset

مقدار، مقدار میانه و مقدار میانگین دقت اعتبارسنجی برای هر پنج الگوریتم به صورت نمودار رسم شده است. فاصله بین حداکثر و حداقل دقت به دست آمده در ۱۰ اجرای مستقل برای الگوریتم پیشنهادی کمتر از چهار الگوریتم دیگر است، این نشان می دهد که نتایج به دست آمده از الگوریتم پیشنهادی قابل اعتماد است؛ بنابراین می توان نتیجه گرفت که الگوریتم پیشنهادی پایدارتر از دیگر الگوریتم ها است.



شکل ۲- عملکرد شاخص ارزیابی در پنج الگوریتم با مجموعه داده



شکل ۳- عملکرد شاخص ارزیابی در پنج الگوریتم با مجموعه داده

۶- نتیجه گیری و کارهای آینده

نتایج آزمایش ها نشان دهنده برتری روش پیشنهادی نسبت به سایر روش های مرز دانش از نظر معیارهای Accuracy, Recall, Precision و F1-score است. در مدل پیشنهادی سعی بر این بود که روشی برای ساختن ارتباط مناسب بین ویژگی های کاربران و تصمیم اعتماد ارائه شود. البته لازم به ذکر است که اعتماد یک موضوع کاملاً ذهنی است و ممکن است افراد در شرایط مختلف تصمیمات مختلفی بگیرند. تا زمانی که ویژگی های افراد بهتر و بیشتر قابل دریافت باشد، تصمیم گیری بهتر و بادقت بیشتری همراه است. باتوجه به ماهیت اعتماد و عدم کارایی مدل کلی و فراگیر برای همه کاربران آنچه مشهود است در نظر گرفتن شرایط فردی و ذهنی کاربران در رابطه با تصمیم اعتماد است و فرق روش پیشنهادی این مقاله نسبت به تحقیقات قبلی این است که سعی شده از ارائه مدل ارزیابی کلی برای همه کاربران پرهیز شود و ارزیابی فردی جایگزین

یک از پنج رویکرد در جدول ۲ (نتایج مجموعه داده های اپینین) و جدول ۳ (نتایج مجموعه داده های ادوگاتو) گزارش شده است.

۴-۵- آزمایش ها و تفسیر نتایج

نتایج ارزیابی رویکرد پیشنهادی، جهت مقایسه بهتر در شکل ۱ و ۲ به صورت نمودار میله ای رسم شده است.

بر اساس نتایج منطبق بر جداول ۲ و ۳، روش پیشنهادی عملکرد مناسب تری نسبت به مدل های مقایسه شده دارد که این بیان کننده صحت و دقت مناسب تر تصمیمات اعتماد است. باتوجه به ترکیب شواهد و اینکه تصمیم اعتماد به صورت مقایسه انجام می پذیرد، می توان برای تصمیم اعتماد، درجه اعتماد را بیشتر از ۰/۶ در نظر گرفت. اگر دو تصمیم اعتماد بر پایه مقادیر ۰/۹ و ۰/۷ باشد مسلم است که تفاوت بین این دو تصمیم وجود دارد و ۰/۹ دقت بیشتری دارد که این نکته را در مدل پیشنهادی با استفاده از معیار اطمینان اعتماد در نظر گرفته می شود.

جدول ۲- مقایسه ارزیابی هر رویکرد به درصد برای مجموعه داده اپینین

رویکرد	Accuracy	Recall	Precision	F1-score
مدل RTARS [۱۸]	۷۲،۰۷	۷۳،۶۹	۷۱،۹۳	۷۲،۷۹
مدل ABC [۲۲]	۷۵،۹۴	۸۰،۳۷	۷۵،۵۳	۷۷،۸۷
DSL-STM [۱۵]	۸۸،۶۷	۹۴،۸۳	۸۸،۳۷	۹۱،۴۳
Automata [۱] Algorithm	۸۹،۱۴	۹۱،۸۷	۸۹،۱۷	۹۰،۴۹
مدل پیشنهادی	۹۲،۴۶	۹۵،۳۲	۹۲،۵۴	۹۳،۹۰

جدول ۳- مقایسه ارزیابی هر رویکرد به درصد برای مجموعه داده ادوگاتو

رویکرد	Accuracy	Recall	Precision	F1-score
مدل RTARS [۱۸]	۶۸،۷۴	۷۰،۲۲	۶۸،۴۶	۶۹،۳۲
مدل ABC [۲۲]	۷۱،۱۶	۷۴،۷۳	۷۱،۴۵	۷۳،۰۵
DSL-STM [۱۵]	۸۵،۷۹	۹۰،۳۵	۸۵،۹۱	۸۸،۰۷
Automata [۱] Algorithm	۶۶،۳۷	۹۰،۶۱	۸۷،۳۸	۸۸،۹۶
مدل پیشنهادی	۸۸،۶۰	۹۱،۴۸	۸۸،۷۴	۹۰،۰۸

برای مقایسه پایداری الگوریتم پیشنهادی نسبت به الگوریتم های دیگر، الگوریتم پیشنهادی و چهار الگوریتم مورد مقایسه، ۱۰ بار به صورت مستقل با تغییر لایه های CV به صورت کاملاً تصادفی (بر روی مجموعه داده های اپینین) اعتبارسنجی و اجرا شده است. در هر اجرا میانگین دقت برای مجموعه داده اعتبارسنجی در جدول ۴، گزارش شده است. در شکل ۴ حداکثر مقدار، حداقل

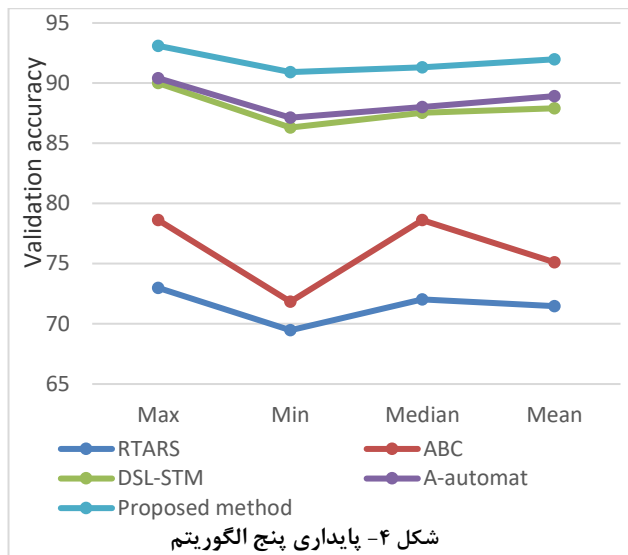
جدول ۴- گزارش دقت مجموعه داده اعتبارسنجی در ۱۰ اجرای مستقل

تکرار	RTARS	ABC	DSL-STM	A-automat	Proposed method
تکرار اول	۷۲،۰۷	۷۵،۹۴	۸۸،۶۷	۸۹،۱۴	۹۲،۴۶
تکرار دوم	۷۰،۳۹	۷۷،۰۶	۸۸،۹۷	۸۹،۰۶	۹۱،۵۸
تکرار سوم	۷۱،۸۲	۷۶،۹۴	۸۶،۶۳	۸۷،۵۲	۹۰،۹۲
تکرار چهارم	۷۱،۵۵	۷۳،۴۵	۸۶،۳۱	۸۷،۱۳	۹۲،۴۲
تکرار پنجم	۷۲،۰۳	۷۸،۶۲	۸۷،۵۳	۸۸،۰۲	۹۱،۳۲
تکرار ششم	۷۲،۹۹	۷۱،۸۴	۹۰،۰۰	۸۸،۴۹	۹۲،۰۱
تکرار هفتم	۷۰،۴۵	۷۳،۶۹	۸۸،۷۱	۹۰،۴۱	۹۱،۰۸
تکرار هشتم	۷۲،۶۵	۷۳،۰۶	۸۶،۷۹	۸۸،۹۴	۹۲،۲۹
تکرار نهم	۷۱،۲۹	۷۵،۴۷	۸۷،۴۶	۹۰،۰۵	۹۳،۱۰
تکرار دهم	۶۹،۴۷	۷۵،۰۸	۸۸،۱۳	۹۰،۳۷	۹۲،۵۹

[۱۱]. بهشید شایسته، وصال حکمی، سید اکبر مصطفوی، احمد اکبری ازیرانی، «ارائه روشی نوین برای محاسبه اعتماد در کاربردهای اینترنت اشیا»، مجله مهندسی برق دانشگاه تبریز، جلد ۵۰، شماره ۲، صفحات ۷۴۳-۷۵۵، ۱۳۹۹.

[۱۲]. مریم فیاض، حامد وحدت نژاد، مهدی خرد، «استنتاج اعتماد در شبکه‌های اجتماعی با ترکیب شبکه عصبی و الگوریتم ژنتیک»، مجله مهندسی برق دانشگاه تبریز، جلد ۵۰، شماره ۱، صفحات ۳۳۱-۳۴۰، ۱۳۹۹.

- [13]. X.Qiu, L.Zhang, S.Wang, "A Trust Transitivity Model Based-on Dempster-Shafer Theory", Journal of Networks, Vol. 5, No. 9, 2010.
- [14]. J.V.S, M.MSK, "Efficient trust management with Bayesian Evidence theorem to secure public key infrastructure-based mobile ad hoc networks", EURASIP Journal on Wireless Communications and Networking, pp.1-27, 2018.
- [15]. W.Abdelghani, I.Amous, C.A.Zayani, F.Sêdes, G.Roman-Jimenez, "Dynamic and scalable multi-level trust management model for Social Internet of Things." The Journal of Supercomputing, 78(6), 8137-8193, 2022.
- [16]. M.Amiri-Zarandi, R.A.Dara, E.Fraser, "LBTM: A lightweight blockchain-based trust management system for social internet of things". The Journal of Supercomputing, pp.1-19, 2022.
- [17]. S.Ma, "Towards Effective Genetic Trust Evaluation in Open Network", IEEE 20th International Conference on High Performance Computing and Communications, 2018.
- [18]. S.Ahmadian, M.Afsharchi, M. Meghdadi, "an effective social recommendation method based on user reputation model and rating profile enhancement." Journal of Information Science, pp. 1-36, 2018.
- [19]. W.Mo, T.Wang, S.Zhang, J.Zhang, "An active and verifiable trust evaluation approach for edge, computing", Journal of Cloud computing Advances, Systems and Applications, pp.1-19, 2020.
- [20]. S.U.Nasir, T.H.Kim, "Trust Computation in Online Social Networks Using Co-Citation and Transpose Trust Propagation", IEEE Access, Vol 8, pp.41362-41371, 2020.
- [21]. Y.Wang, X.Tong, "Trust Prediction Based on Extreme Learning Machine and Asymmetric Tri-Training", IEEE Access, pp.64358-64367, Vol 9, 2021.
- [22]. Sh.Saeidi, "A new model for calculating the maximum trust in Online Social Networks and solving by Artificial Bee Colony algorithm." Computational Social Networks, Vol 7, No. 3, 2020.
- [23]. Z.Wu, L.Tian, Y.Zhang, Z.Wang, "Web User Trust Evaluation: A Novel Approach Using Fuzzy Petri Net and Behavior Analysis", MDPI, Symmetry Journal, No 13, p.1487, 2021.
- [24]. J.Zhang, J.Li, Q.Sun, A.Zhou, "A Comprehensive and Efficient Trust Evaluation Framework for Distributed Networks", Journal of Internet Technology, Vol.19, No.7, 2018.
- [25]. D.Dubois, H.Prade, "Rough fuzzy sets and fuzzy rough sets", International Journal of General System, 17(2-3), 191-209, 1990.
- [26]. LA.Zadeh, "Fuzzy sets". Information Control, PP.338-353, 1965.
- [27]. Z.Pawlak, "Rough sets", International journal of computer & information sciences, 11, pp.341-356, 1982.
- [28]. N.Verbiest, C.Cornelis, F.Herrera, "FRPS: a fuzzy rough prototype selection method", Pattern Recognit, no.10, pp.2770-2782, 2013.
- [29]. Z.G.Liu, J. Dezert, Q.Pan, G.Mercier, "Combination of sources of evidence with different discounting factors based on a new dissimilarity measure". Decis Support Syst NO.52(1), PP.133-141, 2011.
- [30]. Z.Liu, Q.Pan, J.Dezert, G.Mercier, "Credal c-means clustering method based on belief functions". Knowl Based Syst 74, PP.119-132, 2015.
- [31]. Z.G.Liu, Q.Pan, J.Dezert, A. Martin, "Adaptive imputation of missing values for incomplete pattern classification". Pattern Recognit 52:85-95, 2016.



شکل ۴- پایداری پنج الگوریتم

آن شود. از عمده چالش‌های پیشرو می‌توان به استفاده از پایگاه داده به جای مجموعه داده به منظور ارزیابی اعتماد اشاره کرد؛ زیرا باتوجه به معیارهای حال حاضر کاربران و همچنین اطلاعات دقیق‌تر می‌توان نتایج بهتری گرفت و استفاده از پایگاه داده و کاربرد ویژگی‌های بیشتر همواره چشم‌انداز بهتری در ارزیابی اعتماد خواهد داشت و منتج به جامعیت و دقت بیشتر در مبحث اعتماد خواهد بود که این ابزاری مفید و کارا در جهت محاسبات امنیت نرم نیز خواهد بود. البته می‌توان در ارزیابی از روابط ریاضی کامل‌تر منطبق بر مبحث اعتماد استفاده بیشتری کرد که این خود می‌تواند ابعاد جدیدی از کاربرد ریاضی در امنیت نرم را آشکار سازد. در ادامه در زمینه اعتماد کلی و امنیت نرم می‌توان معیاری با عنوان اعتبار کاربر تعریف کرد که این معیار می‌تواند ترکیب قابل‌قبولی از معیارهای شهرت، باور، قابل‌اتکا بودن، اعتماد، اطمینان و کارایی باشد.

مراجع

- [1]. N.Fatehi, H.S.Shahhoseini, J.Wei, C.T.Chang, "An automata algorithm for generating trusted graphs in online social networks." Applied Soft Computing 118 : 108475, 2022.
- [2]. Y.Xu, Z.Feng, X.Zhou, M.Xing, H.Wu, X.Xue, L.Qi, "Attention-based neural networks for trust evaluation in online social networks", Information Sciences, 630, pp.507-522, 2023.
- [3]. N.Fatehi, H.Shahhoseini, "A hybrid algorithm for evaluating trust in online social networks", in: 2020 10th International Conference on Computer and Knowledge Engineering, ICCKE(IEEE), pp.158-162, 2020.
- [4]. H.Shakeri, A.Ghaemi Bafghi, "A Layer Model of a Confidence-aware Trust Management System", International Journal of Information Science and Intelligent System, Vol. 3, pp. 73-90, 2014.
- [5]. P.Alamir, N.Jafari, "Trust Management by Quality of Service & History of Contact in Social Networks", EFTA, Vol.8, pp.53-66, 2013.
- [6]. L.Cholvy, "A Trust Model in the Logical Belief Function Theory", Journal of Applied Logic 13, no. 4, 441-457, 2015.
- [7]. W.Zhang, S.Zhu, J.Tang, N.Xiong, "A Novel Trust Management Scheme Based on Dempster-Shafer Evidence Theory for Malicious Nodes Detection in Wireless Sensor Networks", Springer, Volume 74, Issue 4, pp 1779-1801, 2018.
- [8]. K.Zolfaghar, A.Aghaie, "Evolution Of trust Networks in Social Web Applications Using Supervised Learning", Elsevier Procedia Computer Science, Vol. 3, pp. 833-839, 2011.
- [9]. W.Jiang, G.Wang, J.Wu, "Generating Trusted Graphs for Trust Evaluation in Online Social Networks", Future generation computer systems 31, pp.48-58. Elsevier, 2014.
- [10]. Y.Mahajan, Z.Guo, J.H.Cho, I.R.Chen, "Privacy-Preserving and Diversity-Aware Trust-based Team Formation in Online Social Networks", 2023.