



ORIGINAL RESEARCH PAPER

A trust-based recommender system for e-Learning environment using fuzzy clustering

R. Mohamadrezai, R. Ravanmehr*

Department of Computer Engineering, Central Tehran Branch, Islamic Azad University, Tehran, Iran

ABSTRACT

Received: 09 October 2020
Reviewed: 25 November 2020
Revised: 27 January 2021
Accepted: 27 February 2021

KEYWORDS:

Recommender system
E-Learning
Trust relations
Fuzzy clustering
Weighted association rules

* Corresponding author:

r.ravanmehr@iauctb.ac.ir

① (+9821)44600046

Background and Objectives: Many conventional e-Learning systems are based on static information and consider all learners the same, so they cannot meet their diverse needs and tastes. The main drawback of these systems is ignoring the previous interactions and interests of the learners. The e-learning recommender systems have been introduced with the aim of overcoming these problems and offering the most suitable personalized courses to each learner. The goal of this article is to propose a trust-based e-learning recommender system using fuzzy clustering while taking into account the learners' previous interactions and interests. For this purpose, the weighted association rules and rank prediction were used to produce a candidate list of learning courses and reclassification of the candidate list to generate the final recommendations list.

Methods: In this paper, a novel approach is proposed, which is based on combining the trust relationships among users and their common interests in order to calculate their similarities in an e-Learning recommender system while using fuzzy clustering and weighted association rules, which are aimed at recommending learning courses to the users. In the proposed method, after analyzing the similarities among users and constructing a trust matrix, the next stages are divided into two general phases: the clustering phase of the users and the phase of recommending suitable learning courses for the users. The clustering phase consists of two stages. In the first stage, the optimal number of clusters is obtained using the X-Means algorithm, and in the second stage, the fuzzy C-Means clustering is performed based on the number of clusters obtained. In the recommendation phase for the user, using the weighted association rules and the final clusters obtained for the users, the rank intended by the target user is predicted for each learning item according to the neighbors of the user's cluster. Finally, based on the predicted rankings, N higher ranking course items are suggested as the target user's favorite items.

Findings: Implementation and evaluation of the proposed method on the Moodle dataset demonstrate that with the reduction of the Mean Absolute Error (MAE) and Root Mean Square Error (RMSE), the accuracy of the proposed recommendations is increased, utilizing trust relationships, and the coverage rate of the users and ranks has increased, using fuzzy clustering and weighted association rules, respectively, as compared with the other existing methods. These findings result from employing the fuzzy clustering of users based on their interests and the trust relationships among them, which make it possible for each user to join several clusters with different degrees of membership. Moreover, in utilizing weighted association rules, the association rules that are most compatible with the courses taken by the user are selected. Rules selection scores are calculated on the basis of not only the reliability factors but also a combination of the reliability factors and the user's interest in learning courses.

Conclusions: Utilizing the criterion of trust among users increases the accuracy in choosing neighbors and limits the users' harmful effects and invalid opinions, which will ultimately lead to more accurate recommendations. Also, according to the fuzzy clustering of users, the prediction of the rating of different learning courses is done only based on the neighbors existing in the clusters of the target user. As a result, it will perform more efficiently for the massive volume of information available in an e-Learning system and it shall reduce the problem of data sparsity.



NUMBER OF REFERENCES

50



NUMBER OF FIGURES

3



NUMBER OF TABLES

9

مقاله پژوهشی

سیستم پیشنهاددهنده مبتنی بر اعتماد در محیط یادگیری الکترونیکی با استفاده از خوشه‌بندی فازی

رضوان محمدرضایی، رضا روانمهر*

گروه مهندسی کامپیوتر، دانشکده فنی و مهندسی، واحد تهران مرکزی، دانشگاه آزاد اسلامی، تهران، ایران

چکیده

پیشینه و اهداف: بسیاری از سیستم‌های یادگیری مرسوم مبتنی بر داده‌های ایستا هستند و همه دانش‌آموزان را یکسان و مشابه در نظر می‌گیرند. بنابراین نمی‌توانند پاسخگوی نیازها و سلیق متنوع آن‌ها باشند. مشکل اصلی آن‌ها، در نظر نگرفتن علاقه‌مندی‌ها و تعاملات پیشین کاربران است. سیستم‌های پیشنهاددهنده یادگیری الکترونیکی با هدف غلبه بر این مشکلات و پیشنهاد مناسب‌ترین دوره‌های آموزشی شخصی‌سازی‌شده به هر کاربر مطرح شده‌اند. هدف این مقاله، ارائه یک سیستم پیشنهاددهنده یادگیری الکترونیکی مبتنی بر اعتماد با استفاده از خوشه‌بندی فازی با در نظر گرفتن تعاملات پیشین کاربران و تمایلات آن‌ها است. بدین منظور از کاوش قوانین انجمنی وزن‌دار و پیش‌بینی رتبه برای تولید لیست کاندید دوره‌های آموزشی و رتبه‌بندی مجدد لیست کاندید برای تولید لیست نهایی استفاده شده است.

روش‌ها: در این مقاله یک روش جدید مبتنی بر ترکیب روابط اعتماد بین کاربران و شباهت علایق آن‌ها برای محاسبه میزان تشابه کاربران در یک سیستم پیشنهاددهنده یادگیری الکترونیکی با هدف پیشنهاد دوره‌های آموزشی به کاربران ارائه شده است که از روش خوشه‌بندی فازی و قوانین انجمنی وزن‌دار استفاده می‌کند. در روش پیشنهادی بعد از بررسی شباهت میان کاربران و ساخت ماتریس اعتماد، ادامه مراحل به دو فاز کلی تقسیم می‌شود: فاز خوشه‌بندی کاربران و فاز تولید پیشنهاد دوره‌های آموزشی مناسب برای کاربر. فاز خوشه‌بندی شامل دو مرحله است که در مرحله اول با استفاده از الگوریتم X-Means، تعداد بهینه خوشه‌ها به‌دست می‌آید و در مرحله دوم بر اساس تعداد خوشه‌های به‌دست‌آمده، خوشه‌بندی C-Means فازی انجام می‌شود. در فاز ایجاد پیشنهاد برای کاربر، با استفاده از قوانین انجمنی وزن‌دار و بر اساس خوشه‌های نهایی که برای کاربران حاصل شده‌اند، رتبه موردنظر کاربر هدف، برای هر آیتام آموزشی با توجه به همسایه‌های خوشه‌های کاربر پیش‌بینی می‌شود. در نهایت بر اساس رتبه‌های پیش‌بینی‌شده، N آیتام آموزشی با رتبه بالاتر به عنوان آیتام‌های مورد علاقه کاربر هدف به وی پیشنهاد می‌شوند.

یافته‌ها: پیاده‌سازی و ارزیابی روش پیشنهادی بر روی مجموعه داده Moodle نشان می‌دهد که با کاهش دو معیار میانگین خطای مطلق و خطای جذر میانگین مربعات، دقت پیشنهادی ارائه شده با استفاده از روابط اعتماد افزایش یافته و نرخ پوشش کاربران و رتبه‌ها نیز با استفاده از خوشه‌بندی فازی و قوانین انجمنی وزن‌دار نسبت به روش‌های موجود بهبود یافته است. این نتایج حاصل استفاده از خوشه‌بندی فازی کاربران بر اساس علاقه‌مندی‌های و روابط اعتماد میان آن‌ها است که امکان عضویت هر کاربر را در چند خوشه با درجات عضویت مختلف، قرار می‌دهد. علاوه بر این، در استفاده از قوانین انجمنی وزن‌دار، قوانین انجمنی که بیشترین مطابقت را با دوره‌های گذرانده شده توسط کاربر موردنظر دارند انتخاب می‌شوند. امتیاز‌گذاری انتخاب قوانین، نه تنها بر اساس ضریب اطمینان، بلکه بر اساس ترکیبی از ضریب اطمینان و علاقه‌مندی‌های کاربر به دوره‌های آموزشی، محاسبه شود.

نتیجه‌گیری: بکارگیری معیار اعتماد میان کاربران باعث افزایش دقت در انتخاب همسایه‌ها و محدود کردن اثرات مخرب کاربران و نظرات بی‌اعتبار می‌شود که منجر به ارائه پیشنهادی دقیق‌تری خواهد شد. همچنین با توجه به خوشه‌بندی فازی کاربران، پیش‌بینی رتبه دوره‌های آموزشی مختلف فقط بر اساس همسایه‌های موجود در خوشه‌های کاربر هدف، انجام می‌شود و در نتیجه برای حجم انبوه اطلاعات موجود در یک سیستم یادگیری الکترونیکی، عملکرد کارآتری خواهد داشت و مشکل خلوت بودن داده‌ها را کاهش می‌دهد.

تاریخ دریافت: ۱۸ مهر ۱۳۹۹

تاریخ داوری: ۵ آذر ۱۳۹۹

تاریخ اصلاح: ۰۸ بهمن ۱۳۹۹

تاریخ پذیرش: ۰۹ اسفند ۱۳۹۹

واژگان کلیدی:

سیستم پیشنهاددهنده

یادگیری الکترونیکی

روابط اعتماد

خوشه‌بندی فازی

قوانین انجمنی وزن‌دار

* نویسنده مسئول

✉ r.ravanmehr@iauctb.ac.ir

① ۴۴۶۰۰۰۴۶ - ۰۲۱

مقدمه

سرویس موردنظر که به آن علاقه‌مند است، هدایت می‌کند. منظور از پالایش، انتخاب مناسب‌ترین و مرتبط‌ترین اطلاعات از میان مجموعه داده‌های یک سیستم است. در بین سیستم‌های مختلف دسترسی به اطلاعات، سیستم پیشنهاددهنده نقش مهم و حیاتی برای بهبود

در دهه‌های اخیر افزایش ناگهانی اطلاعات برخط باعث سردرگمی کاربران شده است. سیستم پیشنهاددهنده، یک ابزار پالایش اطلاعات است که کاربر را در یک مسیر شخصی‌سازی‌شده به سمت کالا یا

پیش‌بینی می‌شود نیز افزایش یابد. بنابراین، محاسبه شباهت بین کاربران در صورتی قابل استناد است که آیت‌های رتبه‌بندی شده مشترک زیادی وجود داشته باشد.

در بسیاری از سیستم‌های پیشنهاددهنده به علت نبود یا کمبود رتبه‌های اولیه امکان ارائه پیشنهادها مناسب به کاربر هدف وجود ندارد که به این چالش مشکل شروع سرد گفته می‌شود. بنابراین یکی از مهم‌ترین مشکلات روش پالایش مشارکتی، تشخیص صحیح همسایگی برای کاربر هدف و مشکل شروع سرد است. به طور معمول در سیستم‌های تجاری مانند سیستم‌های یادگیری الکترونیکی، با وجود تعداد زیاد دوره‌ها در زمینه‌های یکسان، آیت‌های رتبه‌بندی‌شده مشترک با تعداد بالا وجود ندارد. از جمله روش‌هایی که می‌توانند به سیستم‌های پیشنهاددهنده مبتنی بر پالایش مشارکتی در انتخاب همسایگی کمک کند، روش‌های مبتنی بر خوشه‌بندی می‌باشند. الگوریتم k-Means یکی از معروف‌ترین و پرکاربردترین روش‌ها برای خوشه‌بندی داده‌ها است که در آن، با تعیین تابع هدف بر اساس میانگین فاصله اعضای هر خوشه نسبت به میانگین آن‌ها، عمل می‌کند و به شکلی داده‌ها را در خوشه‌ها قرار می‌دهد تا میانگین مجموع مربعات فاصله‌ها در خوشه‌ها، کمترین مقدار را داشته باشد. داده‌ها پس از تعدادی تکرار به k خوشه مختلف تقسیم می‌شوند. یکی از مشکلات الگوریتم k-Means و اکثر الگوریتم‌های سنتی خوشه‌بندی این است که برای هر نمونه داده، دو حالت تعلق یا عدم تعلق به خوشه موردنظر در نظر گرفته می‌شود و هر نمونه فقط می‌تواند متعلق به یک خوشه باشد. اما در خوشه‌بندی فازی یک نمونه می‌تواند با درجه عضویت‌های مختلف به چندین خوشه مختلف تعلق داشته باشد که در این حالت نتایج انطباق بیشتری با واقعیت دارند [۶]. با خوشه‌بندی فازی کاربران، پیش‌بینی رتبه آیت‌های مختلف بر اساس همسایه‌های موجود در خوشه کاربر هدف، انجام می‌شود و علاوه بر بهبود مشکل شروع سرد، مقیاس‌پذیری را نیز افزایش می‌دهد.

یکی دیگر از مشکلات سیستم‌های پیشنهاددهنده مبتنی بر پالایش مشارکتی، حملات افراد بدخواه می‌باشد [۷]. یک کاربر بدخواه می‌تواند سیستمی که دارای فرآیند پیشنهادکننده شناخته شده‌ای است و امتیاز دهی کاربرانش نیز مشخص است، به راحتی مورد هجوم قرار داده و یک‌سری آیت‌های نادرست را به کاربر هدف پیشنهاد نماید. در واقع، حمله‌کننده می‌تواند با وارد کردن یک کپی از پروفایل کاربر، به راحتی نتایج سیستم را به نفع خود تغییر دهد. به عنوان مثال ممکن است شخصی در مورد یک دوره آموزشی نامناسب، نظری خوب بنویسد تا ارزش آن را افزایش دهد. اگر با هدف تشخیص جعلی بودن نظرات کاربران، نظرات یک به یک خوانده شوند تا جعلی بودن آن‌ها و مشخص شود، طراحی و ارزیابی الگوریتم‌های شناسایی کار بسیار دشوار و هزینه‌بر است. برای حل این مشکل، سیستم‌های پیشنهاددهنده مبتنی بر اعتماد ارائه شده‌اند که از یک شبکه اعتماد که متشکل از امتیازات مربوط به معتبر بودن کاربران است، برای ارائه پیشنهادها بر اساس نظرات کاربرانی که مورد اعتماد کاربر هدف هستند، استفاده

کسب‌وکار و تسهیل تصمیم‌گیری برای کاربران ایفا می‌کند [۱]. به‌طور کلی، لیست توصیه‌ها بر اساس اولویت و علایق کاربران، ویژگی‌های آیت‌ها، تعاملات گذشته کاربر و برخی اطلاعات اضافی دیگر مانند داده‌های زمانی و مکانی ایجاد می‌شود [۲]. سیستم‌های پیشنهاددهنده به‌طور عمده به پالایش مشارکتی، پالایش مبتنی بر محتوا و ترکیبی بر اساس انواع داده‌های ورودی طبقه‌بندی شده‌اند [۳]. این سیستم‌ها به موازات وب در حال پیشرفت هستند و استفاده از آن‌ها در سال‌های اخیر در اینترنت رشد زیادی داشته است و در حوزه‌های مختلفی نظیر تجارت الکترونیک، یادگیری الکترونیک، پیشنهاد کتاب، موسیقی، فیلم و غیره استفاده می‌شوند. هدف یک سیستم پیشنهاددهنده، ایجاد توصیه‌هایی است که از هر نظر برای کاربران مناسب باشد و به همین جهت این سیستم‌ها سعی در ایجاد توازن بین دقت، تازگی، پراکندگی و ثبات در توصیه‌ها دارند.

یکی از انواع سیستم‌های پیشنهاددهنده، سیستم‌های پیشنهاددهنده در حوزه یادگیری الکترونیکی هستند. یادگیری الکترونیکی به معنای کاربرد تکنولوژی اینترنت برای فراهم کردن امکان آموزش دانش‌آموزان در همه زمان‌ها و در همه مکان‌ها است [۴]. یکی از بارزترین فواید آن در برابر آموزش‌های حضوری، تطبیق‌پذیر بودن بیشتر آموزش‌های الکترونیک است. در واقع، در محیط‌های آموزشی حضوری، یک معلم تمام دانش‌آموزان یک کلاس را به یک شکل راهبری کرده و برای همه آن‌ها از یک استراتژی آموزش استفاده می‌کند. با استفاده از سیستم‌های یادگیری الکترونیک این امکان برای دانش‌آموزان فراهم می‌شود تا استراتژی آموزشی مناسب خود را بر اساس توانایی‌ها و نیازمندی‌هایشان انتخاب کنند که در نتیجه افزایش بهره‌وری آموزشی را در برخواهد داشت. از طرفی، حجم بسیار زیاد اطلاعات مربوط به دوره‌های آموزشی و نیز نداشتن تجربه و تخصص لازم برای انتخاب بهینه، یافتن دوره‌های آموزشی موردنیاز و علاقه دانش‌آموزان را دشوار و زمان‌بر کرده است. سیستم‌های پیشنهاددهنده یادگیری الکترونیکی بر اساس قابلیت‌ها، نیازمندی‌ها و ویژگی‌های شخصی دانش‌آموزان، دوره‌های آموزشی مفید و موثری را به آن‌ها پیشنهاد داده و از اتلاف زمان برای جستجو جلوگیری می‌کنند [۵].

یکی از مهم‌ترین و پرکاربردترین روش‌های سیستم‌های پیشنهاددهنده، پالایش مشارکتی است [۱]. این روش به منظور تولید پیشنهادها مناسب برای کاربر هدف، از نظرات کاربران مشابه با کاربر هدف (به‌عنوان کاربران همسایه) در فرآیند تولید پیشنهاد استفاده می‌کند. با وجود این‌که روش پالایش مشارکتی به صورت گسترده در سیستم‌های پیشنهاددهنده مورد استفاده قرار گرفته است، اما دارای چالش‌هایی نیز هست. سیستم‌های مبتنی بر پالایش مشارکتی برای محاسبه شباهت بین کاربران از امتیازات و رتبه‌هایی که آن‌ها به آیت‌های رتبه‌داده‌شده مشترک می‌دهند، استفاده می‌کنند. اگر یک روش بتواند همسایه‌های مناسبی را برای کاربر هدف تعیین کند، قادر خواهد بود رتبه‌های موردنظر را با دقت بالایی تخمین بزند تا درصد رتبه‌هایی که به درستی

با استفاده از الگوریتم پالایش مشارکتی به پیشنهاد آیتم به کاربران می‌پردازد.

دادو (Dahdouh) و همکاران [۱۱] با استفاده از حجم انبوه رویدادهای ثبت شده کاربران موجود، یک روش با استفاده از فناوری‌های داده‌های کلان بر مبنای چارچوب اسپارک ایجاد کرده‌اند. این روش ابتدا رویدادهای ثبت شده موجود از رفتار کاربران را مورد بررسی قرار داده و سپس به استخراج شباهت میان کاربران با استفاده از قوانین از پیش تعریف شده می‌پردازد.

میلیسویچ (Miličević) و همکاران [۱۲] از ترکیب روش‌های مبتنی بر برجسب و هم‌چنین خوشه‌بندی برای ایجاد پیشنهاد در محیط آموزش دوره‌های برنامه‌نویسی استفاده کرده‌اند که ضمن افزایش دقت پیشنهادها باعث استفاده بهینه‌تر از حافظه در حجم انبوه داده‌ها شده‌اند. وان (Wan) و همکاران [۱۳] یک روش محتوا محور معرفی کرده‌اند که برخلاف کارهای پیشین تمرکز آن بر روی مشخصه‌های آیتم‌های آموزشی می‌باشد. این در حالی است که در روش‌های پیشین تمرکز اصلی سیستم‌ها بر روی یادگیرندگان و تغییرات احتمالی آن‌ها و تحلیل این تغییرات بوده است.

الباتاین (Albatayneh) و همکاران [۱۴] برای انتخاب پیشنهادها مناسب، در بین نظرات افراد در تالارهای گفتگو که پیرامون دوره‌های آموزش الکترونیکی می‌باشند جستجو کرده و سعی کرده‌اند با بررسی محتوای هر نظر و با در نظر گرفتن امتیازهای مثبت/منفی داده شده به آن‌ها اقدام به ارائه نظرهای مناسب در مورد دوره‌های آموزشی نمایند. بنابراین نظرات نامرتبط از این طریق حذف می‌شوند تا کاربر بتواند بر اساس داده‌های درست انتخاب بهتری داشته باشد.

در بررسی سیستم‌اتیکی که به تازگی خانال (Khanal) و همکاران انجام داده‌اند، سیستم‌های پیشنهاددهنده در یادگیری الکترونیکی به چهار دسته مبتنی بر محتوی، مبتنی بر پالایش مشارکتی، مبتنی بر دانش و روش‌های ترکیبی تقسیم شده‌اند [۱۵]. با توجه به یافته‌های این تحقیق، تکنیک‌های یادگیری ماشین، الگوریتم‌ها، مجموعه داده‌ها و ارزیابی خروجی، مؤلفه‌های لازم برای یک سیستم کارآمد پیشنهاددهنده در یادگیری الکترونیکی هستند.

یکی از زمینه‌های مهم در یادگیری الکترونیک، انتخاب طرح‌های یادگیری برای آموزگاران است. در [۱۶] با بررسی الزام پشتیبانی آموزگاران به عنوان طراحان آموزش، سیستمی ارائه شده است که حاصل یکپارچگی سیستم پیشنهاددهنده و سیستم مدیریت فعالیت یادگیری است. سیستم مدیریت فعالیت یادگیری یکی از محیط‌هایی است که آموزگاران از آن برای ایجاد، مدیریت و تصویب طرح‌های یادگیری استفاده می‌کنند. سیستم پیشنهاددهنده ارائه شده بر اساس علاقه-مندی‌ها و نیازهای آموزگاران، از میان طرح‌های یادگیری که توسط آموزگاران دیگر ارائه داده شده‌اند، به معلم‌های جدید طرح‌های آموزشی پیشنهاد می‌دهد.

می‌کنند. این سیستم‌ها با استفاده از اطلاعات پیش‌زمینه‌ای، نظرات جعلی را شناسایی می‌کنند.

پیش از این، سیستم‌های پیشنهاددهنده یادگیری الکترونیکی با رویکردهای مختلف، با هدف بهبود یا کاهش مشکلات فوق، ارائه شده‌اند که برخی از آن‌ها در بخش پیشینه تحقیق به تفکیک عملکرد بررسی می‌شوند. در بخش اهداف تحقیق، ضرورت و اهداف روش ارائه شده در این تحقیق و نوآوری‌های روش پیشنهادی ارائه می‌شود.

پیشینه تحقیق

با توجه به رویکرد ارائه شده در مقاله برای ارائه پیشنهادها موثر در حوزه یادگیری الکترونیکی بر اساس اعتماد بین کاربران با استفاده از خوشه‌بندی فازی، پژوهش‌های مرتبطی که بررسی شده‌اند به سه زیر بخش تقسیم می‌شوند: سیستم‌های پیشنهاددهنده در زمینه یادگیری الکترونیکی، سیستم‌های پیشنهاددهنده مبتنی بر اعتماد، سیستم‌های پیشنهاددهنده مبتنی بر اعتماد در زمینه یادگیری الکترونیکی.

سیستم‌های پیشنهاددهنده در زمینه یادگیری الکترونیکی

با توسعه محیط‌های پیشرفته یادگیری الکترونیکی و افزایش تعداد منابع آموزشی موجود، شخصی‌سازی به دلیل تفاوت‌ها در پس‌زمینه، اهداف، قابلیت‌ها و سبک‌های متفاوت یادگیری دانش‌آموزان، به یک ویژگی مهم در سیستم‌های یادگیری الکترونیکی تبدیل شده است تا به کمک آن بتوان پیشنهادها موثر و کارا و با دقت مناسب ارائه کرد. در ادامه خلاصه‌ای از مهم‌ترین روش‌های ارائه شده برای سیستم‌های پیشنهاددهنده را در محیط‌های یادگیری الکترونیکی خواهیم دید. یکی از اولین تحقیقات اساسی در این زمینه مقاله میلیوویچ (Milicevic) و همکاران [۸] است که برای سیستم پیشنهاددهنده آموزش برنامه‌نویسی پروتوس پیشنهاد داده شده است. این سیستم به صورت خودکار، علاقه‌مندی‌های دانش‌آموزان را با سطح دانش آن‌ها مطابقت داده است. روش پیشنهادی، الگوهای مختلفی از سبک‌های یادگیری و عادت‌ها و رفتارهای دانش‌آموزان را از طریق تست کردن سبک‌های یادگیری دانش‌آموزان و کاوش لیست رویدادهای سرورها شناسایی می‌کند.

مسود (Masud) [۹] یک چارچوب برای یادگیری مشارکتی ارائه می‌دهد که شامل قابلیت همکاری داده‌های معنایی، مدیریت فراداده‌های توزیع شده و پردازش پرس‌وجوهای عامل‌گرا است. البته این سیستم فاقد یک معماری یکپارچه است که قادر به پذیرش داده‌های زمینه‌ای باشد.

تاروس (Tarus) و همکاران [۱۰] یک رویکرد جدید با ترکیب روش‌های الگوکاوی ترتیبی و آگاهی از متن ایجاد نموده‌اند. این روش ابتدا به دنبال کشف شباهت کاربران از طریق سطح دانش و اهداف آن‌ها است و سپس از روش الگوکاوی ترتیبی برای کاوش در میان سیاهه‌های موجود از رفتار کاربران در سیستم استفاده می‌کند. درنهایت با ترکیب این دو مقدار و

در [۲۱] از اطلاعات مربوط به اعتماد میان کاربران برای غلبه بر مشکلات ناشی از روش پالایش مشارکتی استفاده شده است. در روش ارائه شده، وزن‌هایی برای میزان اعتماد میان کاربران که بر اساس معیارهای مختلف محاسبه شده‌اند، در نظر گرفته می‌شود که با هدف ارائه پیشنهادهاى دقیق و شخصی‌سازی‌شده، به صورت تکراری بهینه می‌شود.

در [۲۲] بر اساس پروتکل‌های امنیتی و رمزگذاری همومورفیک، دو روش به نام‌های PPMARS-C و PPMARS-S برای پیشنهاددهی برنامه‌های کاربردی موبایل با حفظ حریم خصوصی ارائه شده است. روش PPMARS-C برای حفظ محرمانگی در محیط‌های سرویس ابری ارائه شده است و PPMARS-S در زمینه شبکه‌های اجتماعی مورد استفاده قرار می‌گیرد.

تعدادی از سیستم‌های پیشنهاددهنده مبتنی بر اعتماد از تکنیک‌های تجزیه ماتریس برای حل مشکلات مدل کردن رفتارهای متوالی کاربران و انتشار اعتماد استفاده می‌کنند در حالی که تمایل‌های ضمنی کاربران را نادیده می‌گیرند. برای پیش‌بینی امتیازات کاربران در جدول کاربر-آیتم در روش ارائه شده در [۲۳] از ترکیب اعتماد، توالی علاقه‌مندی‌ها و تمایل‌های ضمنی کاربران به عنوان ورودی استفاده شده است.

جیانگ (Jiang) و همکاران در مدل خود از الگوریتم بهبودیافته slope one در ساخت پیشنهادها بر مبنای رابطه بین کاربران استفاده کرده‌اند تا با ترکیب روابط اعتماد بین کاربران و شباهت موجود بین آن‌ها باعث بهبود دقت پیشنهادها شوند [۲۴].

پان (Pan) و همکاران [۲۵] یک مدل جدید به نام RoleTS برای ایجاد پیشنهادها بر اساس روابط بین کاربران در محیط یادگیری الکترونیکی ارائه کرده‌اند. اساس کار این روش به این گونه است که ابتدا به استخراج اعتماد نهفته بین کاربران می‌پردازد تا بتواند به صورت ضمنی اعتماد بین آن‌ها را حدس بزند و اقدام به تعریف دو نقش فرد اعتمادکننده و فرد مورداعتماد کند.

سیستم‌های پیشنهاددهنده مبتنی بر اعتماد در زمینه یادگیری الکترونیکی

به دلیل آن‌که استفاده از اعتماد بین کاربران می‌تواند به عملکرد سیستم‌های پیشنهاددهنده کمک شایانی کند، استفاده از این نوع سیستم‌های پیشنهاددهنده برای برطرف کردن مشکلات موجود در زمینه یادگیری الکترونیکی نیز در سال‌های اخیر مورد توجه قرار گرفته است.

دیودی (Dwivedi) و همکاران [۲۶] با استفاده از سطح دانش و سبک یادگیری کاربران و ترکیب آن‌ها اقدام به محاسبه اعتماد میان کاربران در محیط یادگیری الکترونیکی کرده‌اند. طبق این روش اگر سبک یادگیری و سطح دانش دو کاربر به هم دیگر شبیه باشد، اعتماد آن‌ها به یکدیگر بیشتر می‌شود. همچنین کاربران با سطح دانش پایین‌تر به کاربران با سطح دانش بالاتر از خود اعتماد بیشتری در این روش خواهند داشت.

یک سیستم یادگیری الکترونیکی شخصی‌سازی شده و قابل تطبیق در [۱۷] با هدف فراهم کردن یک چارچوب در جهت توسعه محیط‌های یادگیری همه جانبه و برای دانش‌آموزانی که از برنامه‌های مطالعه متداول نمی‌توانند استفاده کنند، ارائه شده است. این سیستم، اطلاعات را از منابع موجود در وب استخراج کرده و با به کارگیری آنتولوژی و پردازش زبان طبیعی بر روی اطلاعات پیش زمینه‌ای یادگیری دانش-آموزان و نیازمندی‌های آن‌ها، پیشنهاد دوره‌های آموزش سازماندهی و ارائه می‌شود.

یک سیستم پیشنهاددهنده ترکیبی به عنوان افزونه‌ای برای سیستم مدیریت یادگیری مودل ارائه شده است که می‌تواند مجموعه‌ای از مخازن یادگیری سازگار با استاندارد مودل را پشتیبانی کند [۱۸]. این سیستم یک لیست رتبه‌بندی از آیتم‌های آموزشی را به دنبال یک سؤال مبتنی بر کلمات کلیدی ساده پیشنهاد می‌کند.

سیستم‌های پیشنهاددهنده یادگیری الکترونیکی تلاش می‌کنند تا موارد آموزشی را به صورت شخصی‌سازی شده به کاربران پیشنهاد دهند. رویکرد ارائه شده در [۱۹] برای کاربران موارد یادگیری شخصی‌سازی شده‌ای را ارائه می‌کند که منطبق با علاقه‌مندی‌ها، سلیقه و دانش پیش‌زمینه‌ای آن‌ها و همچنین ظرفیت حافظه کاربران است. این روش بر اساس ترکیبی از روش پالایش مشارکتی و پالایش مبتنی بر محتوی پیاده‌سازی شده است.

مقاله [۵] بر اساس مشخصات دوره‌های آموزشی و سبک یادگیری دانش‌آموزان پیشنهادهای شخصی‌سازی شده به آن‌ها ارائه می‌دهد. این روش برای ارائه سبک یادگیری موردعلاقه هر کاربر و مشخصات دوره-های آموزشی از مدل فلدر و سیلورمن استفاده می‌کند.

سیستم‌های پیشنهاددهنده مبتنی بر اعتماد

سیستم‌های پیشنهاددهنده مبتنی بر اعتماد بر اساس روابط اعتماد بین آن‌ها، پیشنهادها را به کاربران خود ارائه می‌نمایند. در ادامه خلاصه‌ای از مهم‌ترین سیستم‌های پیشنهاددهنده ارائه شده بر مبنای روابط اعتماد را خواهیم دید.

یک روش پیشنهاددهنده مبتنی بر اعتماد در [۲۰] ارائه شده است که شامل سه گام اصلی است. در گام اول یک معیار شباهت پیشنهاد داده شده که از ترکیب اعتماد و منابع اطلاعاتی غنی مربوط به امتیازات کاربران برای تشخیص شباهت میان کاربران، حذف کاربران غیرمشابه و غلبه بر خلوت بودن داده‌ها استفاده می‌کند. خلوت بودن داده‌ها در سیستم‌های پیشنهاددهنده زمانی اتفاق می‌افتد که تعداد رتبه‌های معلوم از تعداد رتبه‌هایی که برای پیش‌بینی نیاز هستند، بسیار کمتر باشد.

در گام دوم، الگوریتم کلونی مورچه‌ها برای اختصاص وزن به کاربران برای نشان دادن میزان شباهت آن‌ها به کاربر هدف به کار برده می‌شود. در نهایت در گام سوم، از مجموعه‌ای از اطلاعات کاربرانی که بیشترین میزان شباهت به کاربر هدف را داشته‌اند، برای پیش‌بینی امتیازات نامعلوم کاربر هدف به آیتم‌ها استفاده شده است.

جسته‌اند. بر اساس بررسی انجام شده تحقیقات بسیار اندکی بوده است که برای پیشنهاد دوره‌های آموزشی به کاربر هدف همزمان از روابط اعتماد بین کاربران و در نظر گرفتن تمایلات آن‌ها استفاده کرده باشد به گونه‌ای که اثرات مخرب کاربران غیرمرتبط و نظرات بی‌اعتبار در پیشنهاد نهایی کمترین اثرگذاری را داشته باشد. تفاوت میان روش پیشنهادی با دیگر روش‌های مبتنی بر اعتماد، ارائه مدلی مبتنی بر خوشه‌بندی فازی و قوانین انجمنی وزن‌دار در محیط یادگیری الکترونیکی است. بر همین اساس در این مقاله، رویکردی مبتنی بر مدل با استفاده از روش خوشه‌بندی فازی و همچنین روابط اعتماد به منظور افزایش کارایی سیستم‌های پیشنهاددهنده یادگیری الکترونیکی ارائه شده است.

به دلیل عدم قطعیت علاقه‌مندی‌های کاربران، قرار دادن آن‌ها در یک خوشه مشخص نمی‌تواند سلیقه آن‌ها را به درستی مدل کند. بنابراین با در نظر گرفتن این مساله، روش پیشنهادی، با استفاده از روش خوشه‌بندی فازی، کاربران را بر اساس علاقه‌مندی‌های و روابط اعتماد میان آن‌ها، مدل می‌کند تا هر کاربر مبتنی بر ویژگی‌های آن بتواند در چند خوشه با درجات عضویت مختلف عضو شود. با توجه به زمان‌بر بودن انجام محاسبات برای تمامی کاربران، اجرای محاسبات تنها بر روی کاربران هر خوشه صورت می‌گیرد. این روش تا حد زیادی بر مشکل خلوت بودن داده‌ها غلبه کرده، دقت پیش‌بینی‌ها را افزایش داده و مشکل مقیاس‌پذیری آن‌ها برای داده‌های انبوه را نیز بهبود می‌دهد. از طرف دیگر، استفاده از قوانین انجمنی وزن‌دار باعث می‌شود که با پیدا کردن قوانین انجمنی که بیشترین مطابقت را با دوره‌های گذرانده‌شده توسط کاربر موردنظر دارند، امتیاز انتخاب قوانین را نه تنها بر اساس ضریب اطمینان، بلکه بر اساس ترکیبی از ضریب اطمینان و علاقه‌مندی کاربر به دوره‌های آموزشی در نظر گرفته شود.

در این روش پیشنهادی، علاوه بر ماتریس رتبه‌بندی برای کاربران و دوره‌های آموزشی، ماتریس اعتماد میان کاربران نیز ایجاد می‌شود. سیستم پیشنهادی حملاتی از قبیل ایجاد پروفایل‌های کاربری تقلبی به منظور ایجاد پیشنهادهای غیرموثر برای کاربران را از بین می‌برد و اعتماد یک کاربر به کاربر دیگر را مورد بررسی قرار می‌دهد تا پروفایل‌های تقلبی ایجاد شده در سیستم بی‌اثر شوند.

در روش پیشنهادی بعد از بررسی شباهت میان کاربران و ساخت ماتریس اعتماد، ادامه مراحل به دو فاز کلی تقسیم می‌شود، فاز خوشه‌بندی کاربران و فاز تولید پیشنهاد دوره‌های آموزشی مناسب برای کاربر. فاز خوشه‌بندی از دو مرحله تشکیل شده است که در مرحله اول با استفاده از الگوریتم X-Means، تعداد بهینه خوشه‌ها به دست می‌آید، سپس در مرحله دوم بر اساس تعداد خوشه‌های به دست آمده، خوشه‌بندی Fuzzy C-Means انجام می‌شود. در فاز ایجاد پیشنهاد برای کاربر، با استفاده از قوانین انجمنی وزن‌دار و بر اساس خوشه‌های نهایی که برای کاربران به دست آمده‌اند، رتبه موردنظر کاربر هدف، برای هر آیتام آموزشی با توجه به همسایه‌های خوشه کاربر پیش‌بینی می‌شود. در نهایت بر اساس رتبه‌های پیش‌بینی شده، N آیتام آموزشی با رتبه بالاتر به عنوان آیتام‌های مورد علاقه کاربر هدف به وی پیشنهاد می‌شوند.

باشکاران و سانتی (Bhaskaran and Santhi) [۴] با استفاده از روش خوشه‌بندی کاربران در یادگیری الکترونیکی به پیش‌بینی اعتماد در رایانش‌ابری پرداخته‌اند. در این روش، در یک فرایند تکراری، هر گره میزان تعلق خود به گروه‌ها را با میانگین‌گیری از اطلاعات همسایه‌های خود در هر مرحله انجام می‌دهد. همچنین یک میزان حداکثری برای تعداد گروه‌هایی که یک گره می‌تواند در آن‌ها عضو باشد در نظر گرفته می‌شود.

دنگ (Deng) و همکاران [۲۷] یک مدل جدید به نام NCTRS برای محیط‌های یادگیری الکترونیکی ارائه کرده‌اند که به محاسبه اعتماد از طریق روابط موجود در شبکه‌های اجتماعی می‌پردازد. این روش در حقیقت برای غلبه بر مشکل خلوت بودن داده‌ها به ترکیب امتیازات کاربران، بررسی محتوای آیتام‌ها و همچنین اعتماد حاصل از شبکه اجتماعی مابین آن‌ها می‌پردازد.

وان (Wan) و همکاران برای غلبه بر مشکل کمبود داده‌های ارتباطی بین کاربران در محیط‌های یادگیری الکترونیکی (که باعث ناکارآمدی پالایش مشارکتی می‌شود)، یک روش جدید به نام SI-IFL ارائه نموده‌اند که از ترکیب مدل تأثیرگذاری یادگیرنده LIM و الگو کاوی ترتیبی برای ایجاد پیشنهادهای استفاده می‌کند [۲۸]. این روش از ترکیب مدل نفوذ و اثرگذاری دانش‌آموزان بر یکدیگر، استراتژی پیشنهادهای مبتنی بر خودسازماندهی و همچنین کاوش الگوهای ترتیبی تشکیل شده است. در سیستم‌های پیشنهاددهنده‌ای که مبتنی بر پالایش مشارکتی هستند، به دلیل خلوت بودن مجموعه داده‌ها، اطلاعات کافی برای ارزیابی وجود ندارد و پیشنهادهای ارائه شده از دقت کافی برخوردار نیستند. در [۲۹] یک الگوریتم پیشنهاددهنده بر اساس پالایش مشارکتی مبتنی بر آیتام ارائه شده است که بر اساس مجموعه اثرگذار رفتارهای گروهی در یادگیری الکترونیک عمل می‌کند. در این روش امتیاز مربوط به منابع یادگیری علاوه بر رفتار شخصی دانش‌آموزان، به تعاملات و ارتباطات میان دانش‌آموزان گروه و اثرگذاری آن‌ها بر روی یکدیگر نیز بستگی دارد. این روش علاوه بر ترکیب مورد پیشنهاد با k نزدیکترین همسایه، مورد پیشنهاد را با k' همسایه معکوس خود نیز ترکیب می‌کند.

روش ارائه شده در [۳۰] مکانیزم شهرت را با محیط یادگیری الکترونیک ترکیب کرده و یک سیستم پیشنهاددهنده شخصی‌سازی شده ارائه می‌دهد. در واقع کاربران این سیستم، موارد آموزشی را به عنوان پیشنهاد از کاربران دیگر دریافت می‌کنند که مرتبط با محتویات قبلی مطالعه شده آن‌ها باشد. از شهرت در این روش برای بهبود سطح اطمینان و اعتماد اطلاعات پیشنهاد داده شده استفاده می‌شود. در جدول ۱ به بیان مزایا و معایب برخی از کارهای پیشین پرداخته شده است.

اهداف تحقیق

همان‌طور که در فصل قبل بررسی شد تحقیقات مختلفی در زمینه سیستم‌های پیشنهاددهنده یادگیری الکترونیکی انجام شده است که برخی از آن‌ها نیز از روابط اعتماد میان کاربران دوره‌های آموزشی بهره

جدول ۱: مقایسه تعدادی از سیستم‌های پیشنهاددهنده بخش پیشینه تحقیق

Table 1: Comparison of some systems recommending the background section of the research

مرجع Reference	سال انتشار Year	روش Method	مزایا Advantages	معایب Disadvantages
[9]	2016	این روش راهکارهایی برای افزایش قابلیت همکاری مدیریت فراداده‌های توزیع شده، مفاهیم معنایی و یک رویکرد مبتنی بر عامل برای پشتیبانی از تبادل محتویات یادگیری میان سیستم-های آموزش الکترونیکی مختلف ارائه می‌دهد. The proposed method presents solutions for increasing the interoperability of distributed metadata management, semantic concepts, and an agent-based query processing approach to support the exchange of the learning content from different e-learning systems.	با روش خوشه‌بندی سلسله-مراتبی بر مشکل خلوت بودن داده‌ها غلبه می‌کند. The proposed method overcomes the problem of data sparsity by using the hierarchical clustering method.	فرایند محاسبات بسیار زمان‌بر است. The computing process is very time-consuming.
[8]	2011	یک سیستم پیشنهاددهنده یادگیری الکترونیکی ترکیبی بر اساس شناسایی سبک یادگیری، دانش و علاقه مندی‌های کاربر مبتنی بر سیستم پروتوس است که امتیازات توالی‌های پرتکرار و محتویات یادگیری شخصی‌سازی شده را به یادگیرندگان پیشنهاد می‌کند. A hybrid e-learning recommendation system comprised of the learning style identification, knowledge, and preferences of the learners is based on the protus system which suggests ratings of the frequent sequences and the personalised learning content to the learner.	پیشنهادهای ارائه می‌دهد که با دقت بالایی نیازمندی‌های یادگیرندگان و موارد آموزشی را تطبیق می‌دهد و سطح بالایی از پذیرش یادگیرندگان را دارد. Presents recommendations with high accuracy in matching learners' requirements with learning material and, thus, have a higher level of acceptance by the learners.	دارای مشکل شروع سرد است. It has Cold start problem.
[10]	2018	یک رویکرد پیشنهاددهنده ترکیبی، الگوریتم‌های زمینه آگاه، کاوش الگوهای ترتیبی و پالایش مشارکتی را برای پیشنهاد منابع آموزشی به یادگیرنده‌ها باهم ترکیب می‌کند. A hybrid recommender approach combines context-based algorithms, sequential pattern mining, and collaborative filtering in order to recommend learning resources to the learners.	استفاده از رویکرد مبتنی بر زمینه، صحت پیشنهادات را افزایش می‌دهد. Using a context-based approach can improve the accuracy of recommendations	محاسبه شباهت بر اساس محتویات داده‌های حجیم، هزینه محاسبات را افزایش می‌دهد. Calculating the similarity based on the content for the massive volume of data can increase the cost of computations.
[11]	2019	موتور پیشنهاددهی رفتارها و فعالیت‌های گذشته یادگیرندگان در بستر یادگیری آموزش را با استفاده از روش قوانین انجمنی تحلیل کرده و از چارچوب اسپارک و هادوپ برای سهولت پردازش موازی و کاهش هزینه‌های محاسبات استفاده می‌کند. The recommendation engine analyzes the learners' former behavior and activities within the e-learning platform using the method of association rules and Spark and Hadoop Framework to facilitate parallel processing of data and reduce the computation costs.	برای حجم گسترده‌ای از داده‌ها عملکرد خوبی دارد. It works well in a massive volume of data.	برای محیط‌هایی که تعداد کاربران آن‌ها کم است، مناسب نیست. It is not suitable for the environments with a low number of users.
[5]	2019	رویکرد پیشنهادی مبتنی بر سبک یادگیری فلدنر-سیلورمن است که برای ارائه سبک یادگیری دانش آموز و پروفایل موارد آموزشی استفاده می‌شود. این رویکرد نشان می‌دهد که الگوریتم خوشه‌بندی k-means، معیار شباهت کسینوسی و ضریب همبستگی پیرسون برای پیاده‌سازی سیستم‌های پیشنهاد-دهنده موارد آموزشی موثر هستند. The proposed approach is based on the Felder and Silverman learning style which is used to propose both the student learning styles and the learning items profiles. This approach shows that the K-means clustering algorithm, the cosine similarity criterion and the Pearson correlation coefficient are effective tools for implementing learning items recommendation systems.	بر اساس سبک آموزش محبوب کاربران و پروفایل هر کاربر، پیشنهادهای با صحت بالا و نرخ خطای کم ارائه می‌کند. Based on the favorite learning style and profiles of each user, it proposes recommendations with very high accuracy and very low error rates.	دارای مشکل شروع سرد است. It has the problem of cold start.
[19]	2017	یک سیستم پیشنهاددهنده جدید که بر اساس چند معیار شخصی‌سازی شده برای سیستم‌های آموزش الکترونیکی ارائه شده و مبتنی بر ترکیب روش‌های پالایش مبتنی بر محتوا و پالایش مشارکتی است. A new recommendation system is presented on the basis of some personalized criteria for elearning systems and combining collaborative and content-based filtering.	بر اساس دانش زمینه‌ای، علاقه‌مندی‌ها و ظرفیت حافظه دانش‌آموزان، بر مشکل شروع سرد غلبه کرده و پیشنهادات جامعی را ارائه می‌دهد.	در صورت اضافه شدن منابع جدید که توسط یادگیرندگان امتیازدهی نشدند، چالش به وجود می‌آید. الگوریتم مشکل مقیاس‌پذیری و پراکندگی داده دارد. If a new resource is added that is not rated by learners, then it will lead to a challenge. The

		Based on the background knowledge, preferences, and the memory capacity of the students, it overcomes the cold start problem and offers detailed recommendations.	algorithm has the problem of scalability and data scatter.
[16]	2018	این مقاله سیستم Mentor را ارائه می‌دهد که معلمان را برای یافتن طرح‌های آموزشی پیشین، بر اساس علاقه‌مندی‌ها و نیازمندی‌های آن‌ها پشتیبانی می‌کند. Mentor با LAMS که یک ابزار مطرح برای طراحی، مدیریت و ارائه ترتیب‌هایی از فعالیت‌های یادگیری است، یکپارچه شده‌اند. This paper presents the Mentor system, that supports teachers in finding former learning designs on the basis of their needs and preferences. Mentor is integrated into LAMS which is a well-known tool for designing, managing and delivering sequences of learning activities.	با در نظر نگرفتن ارتباطات اجتماعی میان معلمان بخش عظیمی از داده‌های مفید برای ارائه پیشنهاد را نادیده گرفته است. By neglecting social relationships among teachers, much of the useful data for generating recommendations have been overlooked.
		Based on the interests and needs of the teachers and using the learning plans of other teachers, it makes cost-effective and time-saving recommendations for teachers.	
[18]	2020	یک سیستم پیشنهاددهنده ترکیبی به عنوان افزونه برای سیستم مدیریت یادگیری مودل ارائه شده است که می‌تواند مجموعه‌ای از مخازن یادگیری سازگار با استاندارد مودل را پشتیبانی کند. این سیستم یک لیست رتبه‌بندی از آیتم‌های آموزشی را به دنبال یک سؤال مبتنی بر کلمات کلیدی ساده پیشنهاد می‌کند. A hybrid recommendation system is presented as an add-on for the Moodle Learning Management System that can support a collection of learning repositories that are compatible with the Moodle standard. This system suggests a ranking list of educational items following a question based on simple key words.	فراداده‌های تحصیلی مانند کاربرد و مدت زمان پیشنهاد داده شده ممکن است روش را بهبود دهند. در ضمن، روش مشکل شروع سرد دارد. The educational metadata such as the suggested use and duration of content may improve the method. In addition, this method has cold start problem.
[17]	2019	سیستم پیشنهاددهنده تطبیق‌پذیر و شخصی‌سازی شده آموزش الکترونیک یک مدل مبتنی بر آنتولوژی است که از وابستگی‌ها نرخ‌ها و درخت‌های پیمایش برای تولید موارد آموزشی صحیح مطابق با نیازمندی‌های یادگیرندگان استفاده می‌کند. The proposed adaptable and personalized e-learning system is an ontology-based model which uses dependency ratios and parse trees to produce accurate learning items according to the learners' requirements.	کاربران برای استفاده از سیستم به دانش پیش زمینه‌ای و تجربه نیاز دارند. Users need background knowledge and expertise to use this system.
		بر اساس اطلاعات پیش-زمینه‌ای یادگیری کاربران، آنتولوژی و روش‌های پردازش زبان طبیعی، پیشنهادات نسبتاً درستی فراهم می‌شود. On the basis of the students' learning background information, ontology, and natural language processing methods, relatively accurate recommendations are provided.	
[20]	2019	روش پیشنهادی به نام TCFACO از اعتماد میان کاربران به عنوان یک منبع اطلاعاتی غنی و همچنین روش بهینه‌سازی کلونی مورچه‌ها استفاده می‌کند. The proposed method called TCFACO uses both trust statements as a rich source of information and the Ant Colony Optimization (ACO) method.	این روش می‌تواند از طریق فیلتر کردن کاربران، پیش از به کارگیری الگوریتم کلونی مورچه‌ها، فضای حل را کاهش بدهد. The method can reduce the solution space by filtering the users before applying the ant colony algorithm.
[22]	2018	دو طرح پیشنهاددهی برنامه‌های کاربردی موبایل با حفظ حریم خصوصی بر اساس ارزیابی اعتماد پیشنهاد شده است. پیشنهادات بر روی برنامه‌های کاربردی موبایل بر اساس	این روش دارای کارایی بالا، هزینه کم حافظه، مصرف پایین cpu، هزینه‌های

		ارتباطی کم و مصرف باتری پایین است. This method benefits from good efficiency, low memory and communication costs as well as low CPU and battery consumption.	It cannot completely detect the internal malicious users and the external attacks.
[23]	2019	رفتارهای اعتماد کاربران مربوط به نحوه استفاده از برنامه‌های کاربردی ارائه می‌شود. Two projects recommending practical and privacy-preserving cell-phone applications which are based on trust evaluation have been proposed. Recommendations on cellphone practical applications are presented based on the users' trust behaviors related to their usage of cell phones practical applications. روش ISOTrustSeq علاوه بر ارتباطات اجتماعی مبتنی بر اعتماد میان کاربران و رفتارهای ترتیبی آن‌ها، علاقه‌مندی‌های ضمنی کاربران را نیز بررسی می‌کند. این روش مبتنی بر روش تجزیه ماتریس است. In addition to considering social trust-based communications among users and their sequential behaviors, the ISOTrustSeq method also considers their implicit interests. This method is based on the matrix factorization.	داده‌های وسیعی را به عنوان ورودی نیاز دارد و اگر داده‌های کاربر قابل دسترس نباشند روش به شدت تحت تاثیر قرار خواهد گرفت. It requires large amounts of data as the input. If the user data are not accessible, the method will suffer severely.
[4]	2017	یک استراتژی پیشنهاددهی ترکیبی مبتنی بر اعتماد ارائه شده است. این روش از ترکیب الگوریتم کرم شب تاب و K-means برای خوشه‌بندی کاربران بر اساس سبک یادگیری آن‌ها استفاده می‌کند و سپس میانگین وزن‌دهی شده مبتنی بر اعتماد را محاسبه می‌کند. A hybrid trust-based recommendation strategy is presented. It uses a combination of firefly and K-means algorithms for clustering the learners according to their learning styles. Then, the trust-based weighted mean is calculated.	یک روش مبتنی بر اعتماد که سرعت و صحت پیشنهادات را بهبود می‌دهد. A trust-based approach that improves speed and accuracy of recommendations.
[27]	2018	یک مدل جدید پالایش مشارکتی اعتماد آگاه مبتنی بر شبکه عصبی برای پیش بینی امتیازات در محیط یادگیری الکترونیکی به استخراج اطلاعات از چندین منبع (محتوا منابع، امتیازات کاربران و اعتماد اجتماعی) می‌پردازد. این روش شبکه‌های عصبی عمیق و رگرسیون موضوعی مشارکتی را ترکیب کرده و از اعتماد اجتماعی برای پیش‌بینی رتبه‌ها استفاده می‌کند. A novel trust-based- neural network and collaborative filtering model is proposed which exploits information from several sources (content, resources, users' ratings, social trust) to predict the ratings in an e-learning environment. This method combines deep neural network and collaborative topic regression together, and uses social trust for the prediction of ratings..	خلوت بودن داده‌ها را کاهش و توانمندی سیستم را افزایش می‌دهد. با استفاده از روابط چندسطحی در شبکه اجتماعی و یادگیری عمیق پیشنهادهای دقیقی ایجاد می‌کند. It reduces data sparsity and improves the system capability. It also makes accurate recommendations using multilevel relationships in the social networks and deep learning
[26]	2013	یک چارچوب پیشنهاددهی مبتنی بر اعتماد پیشنهاد شده است که منابع یادگیری مورد اعتماد را به یادگیرنده‌ها در محیط‌های آموزش الکترونیکی پیشنهاد می‌دهد. این رویکرد از سبک یادگیری یادگیرنده‌ها و سطح دانش آن‌ها برای استخراج اعتماد میان یادگیرنده‌ها استفاده می‌کند و آن‌ها را با پالایش مشارکتی ترکیب می‌کند. A trust-based recommendation framework is proposed that recommends the reliable learning sources in e-learning environments to the learners. This approach uses both the learners' learning styles and their knowledge levels to elicit trust values among the learners and incorporates them with collaborative filtering.	دارای مشکل شروع سرد و پراکندگی داده است. It has the problem of cold start and sparsity problems.
[28]	2020	یک رویکرد پیشنهاددهی مبتنی بر پالایش ترکیبی پیشنهاد داده شده است که مدل تاثیرگذاری یادگیرنده‌ها، استراتژی پیشنهاددهی خودسازمان ده و کاوش الگوهای ترتیبی را جهت پیشنهاد موارد آموزشی به یادگیرنده‌ها باهم ترکیب می‌کند. A new hybrid filtering recommendation approach is proposed which combines the learners' influence model, a self-organization- based recommendation strategy, and sequential patterns' mining together for recommending the learning items to the learners.	اگر علاقه‌مندی‌های گذشته کاربر موجود نباشند، روش با مشکل جدی مواجه می‌شود. If the former user interests are not available, the method will be faced with a serious challenge.

		خلوت بودن داده‌ها را تا حدی کاهش داده است. By combining the students' tree of interest, and the multi-dimensional features of learning items, and also hidden sequential patterns' mining, it offers acceptable suggestions it almost reduces the problem of data sparsity	
[29]	2017	روش پیشنهادی یک الگوریتم پالایش مشارکتی مبتنی بر مجموعه اثرات رفتارهای گروهی در یادگیری الکترونیکی است. پیشنهادات از ترکیب روش k نزدیک‌ترین همسایه و k نزدیک‌ترین همسایه معکوس ایجاد می‌شوند. The proposed method is a collaborative filtering algorithm based on the influence sets of the collective e-learning behaviors. The recommendations are generated by combining k nearest neighbors and k reverse nearest neighbor.	مقیاس پذیری پایین دارد و به علت ساختار ایستا ماتریس کاربر-آیتم، پیدا کردن همسایه‌ها از طریق اسکن داده‌ها انجام می‌شود که زمان‌بر است. It has low scalability and due to the static structure of the user-item matrix, finding neighbors by scanning data is time-consuming.
[30]	2018	معماری ارائه شده، موارد آموزشی را در سیستم های آموزش الکترونیکی که شهرت کاربرانی که موارد آموزشی را پیشنهاد می‌کنند در آن‌ها مورد بررسی قرار می‌گیرد، پیشنهاد می‌دهد. The presented architecture aims at suggesting the recommendation of learning items in an e-learning environment where the reputation of users who recommend these learning items is considered.	تعداد پیشنهادات ارائه شده و ترتیب قرار گرفتن آن‌ها مناسب نیست. The number of recommendations and their order on the list is not appropriate.

○ استفاده از کاوش قوانین انجمنی وزن دار و پیش‌بینی رتبه برای تولید لیست کاندید دوره‌های آموزشی و رتبه‌بندی مجدد لیست کاندید برای تولید لیست نهایی.

در ادامه مقاله در بخش روش تحقیق، روش پیشنهادی که یک سیستم پیشنهاددهنده مبتنی بر اعتماد با استفاده از خوشه‌بندی فازی می‌باشد، ارائه شده است. در بخش نتایج و بحث، ارزیابی و مقایسه روش پیشنهادی با روش‌های مشابه و در بخش پایانی، نتیجه‌گیری مقاله بیان می‌شود.

روش تحقیق

در این بخش، ابتدا مفاهیم پایه‌ای مرتبط با تکنیک‌های به کار برده شده در روش پیشنهادی که شامل سیستم‌های مبتنی بر اعتماد و خوشه‌بندی هستند، ارائه شده است. سپس در بخش بعدی جزئیات فازهای روش پیشنهادی به تفکیک ارائه خواهد شد.

مبانی اولیه روش پیشنهادی

در این قسمت ابتدا مبانی سیستم‌های پیشنهاددهنده مبتنی بر اعتماد ارائه شده و سپس مبانی اولیه خوشه‌بندی و تئوری فازی مورد بررسی قرار می‌گیرد.

سیستم‌های پیشنهاددهنده مبتنی بر اعتماد

سیستم‌های پیشنهاددهنده مبتنی بر اعتماد، از یک شبکه اعتماد که متشکل از امتیازات مربوط به معتبر بودن کاربران است، برای ارائه

ارزیابی روش پیشنهادی بر اساس مجموعه داده Moodle انجام شده است و نتایج به دست آمده از آزمایش‌ها نشان می‌دهد که روش پیشنهادی در مقایسه با چند روش ارائه شده در این زمینه، خطای جذر میانگین مربعات و میانگین خطای مطلق کمتری به دست آورده و همچنین از مقدار پوشش رتبه بهتری برای کاربران و دوره‌های آموزشی نیز برخوردار است. بر اساس نتایج به دست آمده، دقت پیشنهادها در سیستم پیشنهاددهنده این تحقیق نسبت به روش سنتی پالایش مشارکتی بیشتر شده است. علاوه بر این، به دلیل اینکه در روش پیشنهادی علاوه بر روابط اعتماد، داده‌ها با استفاده از روش خوشه‌بندی فازی مدل می‌شود می‌توان تا حد زیادی بر مشکل خلوت بودن داده‌ها که یکی از مشکلات اساسی و همیشگی سیستم‌های پیشنهاددهنده هست غلبه کرده و دقت پیش‌بینی‌ها افزایش یابد. از طرفی با توجه به زمان‌بر بودن انجام محاسبات برای تمامی کاربران، از طریق خوشه‌بندی و اجرای محاسبات تنها بر روی کاربران هر خوشه، می‌توان با بهبود عملکرد سیستم مشکل مقیاس‌پذیری آن‌ها برای داده‌های انبوه را نیز بهبود داد.

به طور خلاصه می‌توان نوآوری‌های این پژوهش را به ترتیب زیر خلاصه کرد:

- ارائه یک سیستم پیشنهاددهنده یادگیری الکترونیکی مبتنی بر اعتماد با استفاده از خوشه‌بندی فازی و قوانین انجمنی وزن دار با در نظر گرفتن تعاملات پیشین کاربران و تمایلات آن‌ها.
- محاسبه شباهت میان کاربران بر اساس ترکیب روابط اعتماد و معیار ضریب همبستگی پیرسون.

شباهت برای ارائه پیشنهاد استفاده نمی‌شود، در مقابل این نوع از حملات نفوذناپذیر می‌شوند.

سیستم‌های پیشنهاددهنده مبتنی بر اعتماد بر اساس نحوه محاسبه رتبه‌های اعتماد کاربران، به دودسته سیستم‌های مبتنی بر اعتماد صریح و سیستم‌های مبتنی بر اعتماد ضمنی تقسیم می‌شوند. در سیستم‌های پیشنهاددهنده مبتنی بر اعتماد صریح از رتبه‌های اعتمادی که کاربران نسبت به یکدیگر ابراز کرده‌اند، برای تولید پیشنهاد استفاده می‌شود. این سیستم‌ها به دو دسته تقسیم می‌شوند: دسته اول، سیستم‌های پیشنهاددهنده مبتنی بر اعتماد که تنها از رتبه‌های اعتماد کاربران استفاده می‌کنند [۳۶]. دسته دوم، سیستم‌های مبتنی بر اعتماد و عدم اعتماد که از رتبه‌های عدم اعتماد نیز در کنار رتبه‌های اعتماد، در فرایند پیش‌بینی رتبه استفاده می‌کنند [۳۷، ۳۸].

در سیستم‌های پیشنهاددهنده مبتنی بر اعتماد ضمنی، رتبه‌های اعتماد کاربران بدون در اختیار داشتن رتبه‌های اعتماد صریح و بر اساس روش‌های مبتنی بر دانش از تعاملات بین کاربران استنتاج شده و از آن‌ها در فرایند تولید پیشنهاد استفاده می‌شود. این سیستم‌ها بر اساس مدلی که برای محاسبه اعتماد استفاده می‌کنند، به دو دسته تقسیم می‌شوند: دسته اول سیستم‌های مبتنی بر اعتماد محلی هستند که رتبه‌های اعتماد را به صورت شخصی و مستقل برای هر کاربر محاسبه می‌کنند. دسته دوم سیستم‌های مبتنی بر اعتماد سراسری هستند که رتبه اعتماد کاربران را بر اساس شهرت آن‌ها در سیستم، محاسبه می‌کنند [۳۹].

لازم به ذکر است که سیستم توصیه‌گر پیشنهادی در این مقاله در دسته سیستم‌های پیشنهاددهنده مبتنی بر اعتماد ضمنی قرار می‌گیرد.

مفاهیم اولیه خوشه‌بندی و تئوری فازی

خوشه‌بندی با فراهم کردن قابلیت ورود به فضای داده و تشخیص ساختار آن‌ها، یکی از ایده‌آل‌ترین روش‌ها برای کار با دنیای عظیم داده‌ها است [۴۰]. خوشه‌بندی یکی از شاخه‌های یادگیری بدون نظارت است و فرایند خودکاری است که در طی آن، نمونه‌ها به دسته‌هایی که اعضای آن مشابه یکدیگر می‌باشند تقسیم می‌شوند که به این دسته‌ها خوشه گفته می‌شود. بنابراین خوشه مجموعه‌ای از اشیاء است که در آن اشیاء با یکدیگر مشابه بوده و با اشیاء موجود در خوشه‌های دیگر غیرمشابه می‌باشند [۶]. برای تشابه می‌توان معیارهای مختلفی را در نظر گرفت، به عنوان مثال می‌توان معیار فاصله را برای خوشه‌بندی مورد استفاده قرارداد و اشیائی را که به یکدیگر نزدیک‌تر هستند را به عنوان یک خوشه در نظر گرفت که به این نوع خوشه‌بندی، خوشه‌بندی مبتنی بر فاصله نیز گفته می‌شود. روش‌های موجود خوشه‌بندی را می‌توان در پنج دسته کلی روش‌های سلسله‌مراتبی، جزءبندی، مبتنی بر تراکم، مبتنی بر شبکه و مبتنی بر مدل دسته‌بندی کرد [۴۱].

در میان روش‌های مختلف در خوشه‌بندی داده‌ها، روش خوشه‌بندی K-Means و C-Means فازی به‌طور گسترده در زمینه‌های مختلف مورد

پیشنهادها بر اساس نظرات کاربرانی که مورد اعتماد کاربر هدف هستند، استفاده می‌کنند. اعتماد، درواقع میزان باور کاربران نسبت به یکدیگر است که بر اساس معیارهایی همچون توانایی، قدرت و خوبی افراد شکل می‌گیرد. رابطه اعتماد، در اکثر مواقع به صورت یک رابطه یک‌طرفه بیان می‌گردد که در آن به کاربری که به سایر کاربران اعتماد می‌کند، اعتماد-کننده و فردی که درواقع مقصد رابطه اعتماد است، مورداعتماد، گفته می‌شود [۷]. راه‌حل مشترکی که برای چالش مدیریت اعتماد مطرح شده است، ایجاد شبکه اعتماد است. شبکه اعتماد یک گراف جهت‌دار است که گره‌های آن نشان‌دهنده کاربران سیستم و یال‌های آن بر اساس درجه اعتماد کاربران به یکدیگر وزن‌دهی می‌شوند. با انجام این کار، یک شخص می‌تواند به دیگران برای یافتن کاربران مورد اعتماد که پیش ازاین با آن‌ها تعامل نداشته‌اند، کمک کند. تحقیقاتی که تاکنون در زمینه مدیریت اعتماد در شبکه‌های اعتماد انجام شده‌اند به دو دسته تقسیم می‌شوند. دسته اول، روش‌هایی هستند که با استفاده از انتشار اعتماد در شبکه، میزان اعتماد کاربران نسبت به یکدیگر را نتیجه می‌گیرند. اما دسته دوم، با استفاده از روش‌های یادگیری ماشین و تعاملات قبلی کاربران، میزان اعتماد بین آن‌ها را پیش‌بینی می‌کنند.

هرچند که سیستم‌های پیشنهاددهنده مبتنی بر پالایش مشارکتی از نظر معیار دقت، عملکرد خوبی دارند، اما با چالش‌هایی مواجه هستند که سیستم‌های پیشنهاددهنده مبتنی بر اعتماد این مشکلات را برطرف می‌کنند [۳۱، ۳۲]. همان‌طور که در بخش پیشین اشاره شد، سیستم‌های پیشنهاددهنده مبتنی بر پالایش مشارکتی از دو ضعف اساسی در فرایند یافتن کاربران مشابه و حملات افراد بدخواه رنج می‌برند. به طور معمول در یک سیستم پیشنهاددهنده، تعداد آیت‌های دارای رتبه بسیار محدود هستند و به علت پراکندگی داده‌ها، احتمال اینکه دو کاربری که به‌صورت تصادفی انتخاب شده‌اند، به آیت‌های مشترکی رتبه داده باشند، بسیار ضعیف است. این مشکل به خصوص برای کاربران جدید سیستم، شدت بیشتری دارد. برای غلبه بر این مشکل، سیستم‌های پیشنهاددهنده مبتنی بر اعتماد از قابلیت انتشار رتبه‌های اعتماد استفاده می‌کنند و برای تولید پیشنهاد، به دنبال ایجاد یک همسایگی از شبکه‌ی اعتماد کاربر هدف هستند.

یکی دیگر از مشکلات سیستم‌های پیشنهاددهنده پالایش مشارکتی، مشکل حملات افراد بدخواه است. روش‌های مختلفی برای تشخیص و مقابله با حملات به سیستم‌های پیشنهاددهنده معرفی شده است [۳۳، ۳۴، ۳۵]. روش‌های مبتنی بر پالایش گروهی، برای یافتن کاربران مشابه با کاربر هدف از معیارهای شباهت بین کاربران استفاده می‌کنند، که این معیارها بر اساس رتبه‌هایی است که کاربران به آیت‌های مشترک داده‌اند. حمله‌کنندگان سعی می‌کنند با درج پروفایل‌های تقلبی و تخصیص رتبه‌های از پیش تعریف‌شده در این پروفایل‌ها، شباهت خود را با این کاربران به‌منظور پیشنهاد آیت‌های مورد نظر، افزایش دهند. اما در سیستم‌های پیشنهاددهنده مبتنی بر اعتماد، از آنجاکه فقط از معیار

اضافه می‌کند تا به حد بالا برسد. در طول این فرآیند مراکز برای دستیابی به بهترین امتیاز تغییر می‌کنند. روش پیشنهادی در این مقاله، از روش خوشه‌بندی فازی برای مدل کردن کاربران بر اساس روابط اعتماد میان آن‌ها استفاده می‌کند. در واقع فاز خوشه‌بندی رویکرد ارائه شده از دو مرحله تشکیل شده است که در مرحله اول با استفاده از الگوریتم X-Means، تعداد بهینه خوشه‌ها به آمده، سپس در مرحله دوم بر اساس تعداد خوشه‌های به‌دست‌آمده، خوشه‌بندی C-Means فازی انجام می‌شود.

نمای کلی روش پیشنهادی

همان‌طور که در قسمت‌های قبل توضیح داده شد، هدف این مقاله معرفی سیستم پیشنهاددهنده یادگیری الکترونیک مبتنی بر اعتماد ضمنی است که بر اساس میزان اعتماد میان کاربران و علایق آن‌ها، با استفاده از روش‌های خوشه‌بندی فازی و قوانین انجمنی وزن‌دار، توصیه‌های دقیق و موثری تولید می‌کند. روش پیشنهادی این مقاله، مبتنی بر مدل بوده و بر اساس پالایش مشارکتی عمل می‌کند. از طرفی در این روش، برای غلبه بر چالش‌های پالایش مشارکتی و افزایش کارایی سیستم پیشنهاددهنده، از روابط اعتماد بین کاربران نیز استفاده شده است.

در روش پیشنهادی ابتدا دو ماتریس رتبه و اعتماد به‌عنوان ورودی‌های سیستم ایجاد می‌شوند که اولی رتبه کاربران به دوره‌های آموزشی موجود و دومی میزان اعتماد کاربران به یکدیگر را نشان می‌دهد. سپس پنج مرحله (۱) محاسبه شباهت بین کاربران، (۲) تعیین تعداد بهینه خوشه‌ها، (۳) خوشه‌بندی فازی کاربران و کاوش قوانین انجمنی وزن‌دار، (۴) انتخاب دوره‌های آموزشی کاندید بر اساس قوانین انجمنی و پیش‌بینی رتبه و (۵) پیشنهاد دوره‌های آموزشی به کاربر، به ترتیب اجرا می‌شوند. در واقع، ابتدا مقادیر شباهت بین هر جفت از کاربران را با ترکیب روابط اعتماد و شباهت به دست آورده و بعد از آن ادامه فرایند به دو فاز کلی خوشه‌بندی کاربران (مراحل ۲ و ۳) و پیشنهاد دوره‌های آموزشی برای کاربر (مراحل ۴ و ۵) تقسیم می‌شود.

فاز خوشه‌بندی از دو مرحله تشکیل شده است که در مرحله اول با استفاده از الگوریتم X-Means، تعداد بهینه خوشه‌ها به دست آمده، سپس در مرحله دوم این فاز بر اساس تعداد خوشه‌های به‌دست‌آمده، خوشه‌بندی C-Means فازی انجام می‌شود. در فاز دوم روش پیشنهادی، بر اساس خوشه‌های نهایی که برای کاربران به‌دست‌آمده‌اند، رتبه موردنظر کاربر هدف برای هر آیتام آموزشی با توجه به همسایه‌های خوشه کاربر پیش‌بینی می‌شود. در نهایت براساس امتیازهای پیش‌بینی شده، N آیتام به رتبه بالاتر به عنوان آیتام‌های مورد علاقه کاربر هدف به وی پیشنهاد می‌شوند.

در روش پیشنهادی از اعتماد میان کاربران استفاده و اقدام به ایجاد یک شبکه اعتماد برای کاربران شده است، زیرا تحقیقات بسیار زیادی نشان داده است که دقت پیشنهادها در سیستم‌های پیشنهاددهنده مبتنی بر

استفاده قرار می‌گیرد. الگوریتم K-Means یک از معروف‌ترین روش‌ها برای خوشه‌بندی داده‌ها به شمار می‌رود. در این الگوریتم، داده‌ها پس از تعدادی تکرار به k خوشه مختلف دسته‌بندی می‌شوند. یکی از مشکلات الگوریتم K-Means و اکثر الگوریتم‌های سنتی خوشه‌بندی این است که تعلق داده به هر خوشه با عدد صفر و یک مشخص می‌شود. به عبارت دیگر در خوشه‌بندی کلاسیک، هر نمونه ورودی متعلق به یک و فقط یک خوشه است و نمی‌تواند عضو دو خوشه و یا بیشتر باشد. درواقع می‌توان گفت، خوشه‌ها همپوشانی ندارند.

برای حل چالش‌های روش خوشه‌بندی K-Means، الگوریتم خوشه‌بندی فازی ارائه شده است [۴۲]. در خوشه‌بندی فازی یک نمونه می‌تواند متعلق به بیش از یک خوشه باشد. در الگوریتم خوشه‌بندی فازی تعلق هر داده به یک خوشه خاص با یک عدد حقیقی بین صفر و یک مشخص می‌شود. ایده بنیادین در خوشه‌بندی فازی به این ترتیب است که فرض کنیم هر خوشه مجموعه‌ای از عناصر است، سپس با تغییر در تعریف عضویت عناصر در این مجموعه از حالتی که یک عنصر فقط بتواند عضو یک خوشه باشد (حالت افزایی)، به حالتی که هر عنصر می‌تواند با درجه عضویت‌های مختلف به چندین خوشه مختلف تعلق داشته باشد، خوشه‌بندی‌هایی که انطباق بیشتری با واقعیت دارند ارائه کنیم [۴۲]. در سال‌های اخیر نسخه‌های بهبودیافته‌ای از این الگوریتم نیز ارائه شده است. برای حل مشکلات ناشی از مقداردهی اولیه خوشه‌ها، استتکو (Stetco) و همکارانش یک طرح جدید در مقداردهی اولیه خوشه‌ها در خوشه‌بندی فازی C-Means ارائه دادند [۴۳]. روش فازی C-Means به‌طور گسترده در زمینه‌های مختلف مانند سنجش از راه دور، خوشه‌بندی سری‌های زمانی و قطعه‌بندی تصاویر رنگی و ... مورد استفاده قرار می‌گیرد.

یکی دیگر از مشکلات خوشه‌بندی، تعیین تعداد بهینه خوشه‌ها است. یکی از راه‌حل‌ها استفاده از معیارهای ارزیابی مختلفی برای ارزیابی خوشه‌بندی غیرفازی است که از جمله آن‌ها می‌توان به معیار DB، DI، CS، DB اشاره کرد. برای تمام این معیارها، مقدار بیشینه و یا کمینه آن‌ها نشان‌دهنده خوشه‌بندی بهینه مجموعه الگوها و یا داده‌هاست. راه حل دیگر این است که برای حل این مشکل و اجرای حدس اولیه تعیین تعداد خوشه‌ها پیش از اجرای الگوریتم، از الگوریتم X-Means که توسط پلگ و موور (Pelleg and Moore) [۴۴] پیشنهاد داده شده است، استفاده شود. X-Means برای تعیین تعداد خوشه‌ها به‌طور خودکار بر اساس امتیازهای معیار اطلاعات بیزین مطابق با رابطه (۱) عمل می‌کند.

$$BIC(\theta) = L(D) - \frac{1}{2} p \log N \quad (1)$$

در رابطه (۱) $L(D)$ درست‌نمایی مبتنی بر لگاریتم مجموعه داده D بر اساس مدل θ است که تعداد خوشه‌ها را مشخص می‌کند، p تعداد پارامترهای آزاد در مدل و N اندازه مجموعه داده است. این الگوریتم در ابتدا با مقداری برابر K برای حد پایین تعداد مراکز شروع شده و الگوریتم K-Means را بر روی مجموعه داده اجرا می‌کند. سپس هر خوشه والد را به دو خوشه فرزند تقسیم کرده و بر اساس نتایج رابطه (۱) مراکز را

تعداد کمی از دوره‌های آموزشی رتبه داده باشند. بنابراین بسیاری از درایه‌های این ماتریس خالی خواهد بود. هر عنصر در فضای مجموعه U می‌تواند توسط پروفایل کاربران با مشخصه‌های متفاوتی از قبیل شناسه کاربری، سن، جنسیت و غیره مشخص شود. در ساده‌ترین حالت، پروفایل کاربران فقط شامل یک عنصر منحصر به فرد به نام شناسه کاربری است. به‌طور مشابه در فضای مجموعه E ، هر آیت می‌تواند توسط مشخصه‌های آن تعریف گردد. به عنوان مثال در یک سیستم پیشنهاددهنده یادگیری الکترونیکی، هر دوره آموزشی می‌تواند توسط ویژگی‌هایی از قبیل شناسه منحصر به فرد آیت آموزشی، عنوان دوره آموزشی، نوع دوره آموزشی و غیره مشخص شود. معمولاً سیستم‌های پیشنهاددهنده در ابتدای کار از کاربران می‌خواهند درجه علاقه‌مندی خود به هر یک از دوره‌های آموزشی را به وسیله یک رتبه مشخص کنند. این امتیازها با استفاده از ماتریس کاربر-آیت نمایش داده می‌شوند که می‌توان از این ماتریس برای محاسبه شباهت بین کاربران استفاده کرد. در جدول ۲ نمونه‌ای از چند سطر یک ماتریس رتبه آورده شده است و همان‌طور که مشخص هست هر سطر دربرگیرنده یک کاربر و هر ستون یک آیت است.

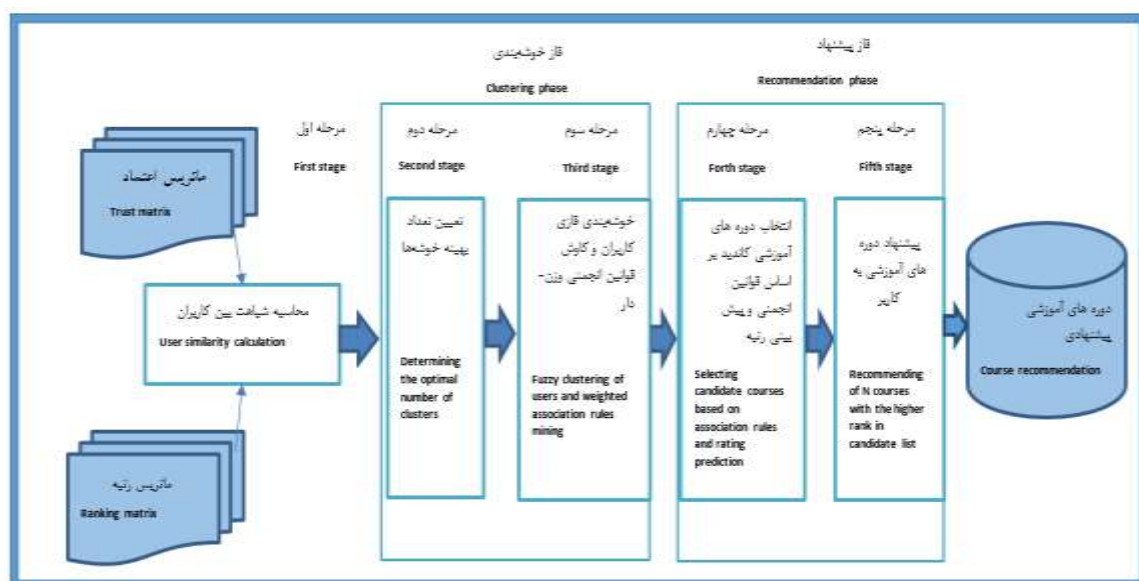
بنابراین اگر تعداد کاربران را n و تعداد دوره‌های آموزشی را m در نظر بگیریم، اولین ورودی سیستم ما یک ماتریس $[n \times m]$ به‌عنوان ماتریس رتبه است. در این مقاله از مجموعه داده Moodle استفاده شده است که در آن کاربران می‌توانند در یک دوره آموزشی از آیت‌های ارائه‌شده مواردی را انتخاب و در آن‌ها شرکت نمایند (در بخش نتایج و بحث که مربوط به شبیه‌سازی و ارزیابی روش پیشنهادی است، به‌طور کامل مجموعه داده Moodle را معرفی خواهیم کرد).

اعتماد نسبت به روش سنتی پالایش مشارکتی بیشتر است [۴۵، ۴۶، ۴۷] همچنین سیستم پیشنهادی اثرات حملاتی از قبیل ایجاد پروفایل‌های کاربری تقلبی به‌منظور ایجاد پیشنهادهای تکراری و غیرموثر برای کاربران را از بین می‌برد، زیرا تنها به دلیل شباهت انتخاب‌های کاربران به یکدیگر پیشنهاد ایجاد نمی‌شود و می‌بایست کاربر به کاربر دیگر اعتماد داشته باشد و این خود باعث می‌شود پروفایل‌های تقلبی ایجاد شده در سیستم بی‌اثر شوند [۴۸].

شکل ۱ چارچوب روش پیشنهادی را به‌طور کامل نشان می‌دهد که در ادامه هر یک از قسمت‌های آن به صورت جداگانه توضیح داده شده و در انتها، شبهه کد روش پیشنهادی بررسی خواهد شد.

ایجاد ماتریس رتبه

در روش پیشنهادی برای محاسبه شباهت بین کاربران باید ماتریس رتبه یا همان ماتریس کاربر-آیت (دوره آموزشی) را به‌عنوان ورودی ایجاد کرده تا با توجه به روابط شباهت، که در اینجا بر اساس معیار ضریب همبستگی پیرسون است، میزان شباهت بین کاربران محاسبه شود. در یک سیستم پیشنهاددهنده با n کاربر و m آیت، مجموعه کاربران به صورت $U = [u_1, u_2, \dots, u_n]$ و مجموعه آیت‌ها نیز به صورت $E = [e_1, e_2, \dots, e_m]$ نشان داده می‌شوند. در واقع کاربران، دانشجویان (دانش‌آموزان) سیستم و آیت‌ها، دوره‌های آموزشی مختلف هستند. ماتریس کاربر-آیت از ورودی‌های اصلی یک سیستم پیشنهاددهنده است. این ماتریس شامل رتبه کاربران به دوره‌های آموزشی موجود در سیستم است. به عبارت دیگر هر سطر نشان‌دهنده یک کاربر و هر ستون یک دوره آموزشی است. ممکن است کاربران به



شکل ۱: چارچوب روش پیشنهادی

Fig. 1: Framework of the proposed method

سطر اول برای ستون سوم مقدار اعتماد برابر ۱ باشد به این معنی است که کاربر سطر اول به کاربر ستون سوم اعتماد دارد. در سیستم‌های پیشنهاددهنده معمولاً تعداد کاربران به حدی زیاد هست که ماتریس اعتماد دارای درایه‌های خالی زیادی است و تعداد کاربر کمی اعتماد به یکدیگر را تایید نموده‌اند. اعتماد در سیستم‌های واقعی در مواردی به صورت صریح و غیرضمنی اندازه‌گیری می‌شوند به این صورت که از کاربر به طور مستقیم در مورد رفتار یا نظرات کاربر دیگر سؤال نمی‌شود. در این روش مبنای اعتماد، اعتباراتی است که کاربران به دست آورده‌اند. در یک سیستم آموزشی معمولاً دانشجویان به دانشجویان قدیمی‌تر از خود که نمرات بالاتری را کسب کرده‌اند برای انتخاب دوره آموزشی اعتماد دارند. در دیتاست مورد استفاده در این مقاله جدول badges در مجموعه داده Moodle نشان‌دهنده این مورد است که کدام کاربران توانسته‌اند با اتمام موفقیت‌آمیز دوره‌های آموزشی و دریافت نمره عالی، مدال‌هایی را کسب کنند و بنابراین می‌توان از آن برای ساخت ماتریس اعتماد استفاده کرد. این جدول شامل فیلدهای نام کاربر، شناسه مدال‌های کسب‌شده و تاریخ دریافت مدال است. به عنوان مثال اگر کاربر اول طبق اطلاعات جدول badges دارای مدال‌های کمتر و یا برابری نسبت به کاربر دوم باشد، کاربر اول به کاربر دوم اعتماد دارد و مقدار درایه ماتریس اعتماد برای کاربر اول و دوم در سطر اول و ستون دوم برابر ۱ مقداردهی خواهد شد. طبق توضیحات بالا می‌توان تعداد مدال‌های هر کاربر را مشخص کرده و اقدام به ایجاد ماتریس اعتماد به عنوان دومین ورودی سیستم نمود و از آن در ادامه مراحل برای ساخت شبکه اعتماد کاربر هدف و همچنین تخمین نهایی شباهت بین هر جفت کاربر استفاده کرد. در نهایت ماتریس اعتماد ساخته‌شده به صورت جدول ۳ خواهد بود.

در روش پیشنهادی برای ایجاد ماتریس رتبه، جدول grades در مجموعه داده ذکر شده به کار رفته است که دربرگیرنده اطلاعات مشارکت کاربران در دوره‌های آموزشی است و شامل فیلدهایی مانند شناسه کاربران، شناسه آیت، علاقه کاربر به آیت، نمره کاربر، زمان مشارکت و غیره می‌باشد. بنابراین بر اساس جدول grades و با استفاده از فیلدهای شناسه کاربر، شناسه دوره آموزشی و میزان علاقه کاربر به دوره آموزشی اقدام به ساخت ماتریس رتبه به عنوان اولین ورودی سیستم می‌شود. با داشتن ماتریس رتبه می‌توان با استفاده از معیار ضریب همبستگی پیرسون که در ادامه شرح داده خواهد شد، اقدام به محاسبه شباهت میان هر جفت کاربر نمود اما برای محاسبه اعتماد میان کاربران نیاز به داشتن ماتریس اعتماد به عنوان ورودی نیز خواهد بود.

ایجاد ماتریس اعتماد

در روش پیشنهادی برای ایجاد شبکه اعتماد به ماتریس اعتماد کاربران نیاز است. برخلاف ماتریس رتبه که مربوط به ارتباط میان کاربران و دوره‌های آموزشی است، در ماتریس اعتماد، ارتباط بین کاربران با یکدیگر نشان داده می‌شود. بنابراین سطر و ستون‌های این ماتریس را مجموعه کاربران تشکیل می‌دهند. اگر تعداد کاربران برابر n در نظر گرفته شود، ماتریس اعتماد یک ماتریس $[n \times n]$ است. معمولاً در سیستم‌های پیشنهاددهنده از کاربران در مورد نظرات کاربران دیگر سؤال می‌شود و اگر کاربر نظر کاربر دیگری را تایید کند درواقع به آن کاربر اعتماد دارد که می‌تواند مقدار ۰ یا ۱ به معنای اعتماد نداشتن یا داشتن و یا مقدار عددی بازه‌ای مثل $[0, 5]$ باشد که میزان اعتماد را دقیق‌تر مشخص نماید. لذا از طریق این ماتریس می‌توان دریافت که هر کاربر به کدام یک از کاربران دیگر اعتماد دارد، برای مثال زمانی که در

جدول ۲: نمونه‌ای از ماتریس رتبه (ماتریس کاربر-آیت)
Table 2: Example of ranking matrix (user-item matrix)

Users کاربرها	Items آیت‌ها							
	I_1	I_2	I_3	I_4	I_5	I_6	I_7	I_8
U_1	4	3	2	4	5	1	5	3
U_2	0	4	3	0	1	3	2	0
U_3	5	1	0	1	0	4	3	0
U_4	1	5	4	0	3	2	5	4
U_5	3	4	5	5	1	0	2	3

جدول ۳: نمونه‌ای از ماتریس اعتماد (ماتریس کاربر-کاربر)
Table 3: An example of a trust matrix (user-user matrix)

Users کاربرها	Users کاربرها							
	U_1	U_2	U_3	U_4	U_5	U_6	U_7	U_8
U_1	0	3	2	4	5	1	5	3
U_2	0	0	3	0	1	3	2	0
U_3	5	2	0	1	0	4	3	0
U_4	1	5	4	0	3	2	5	4
U_5	3	4	5	5	0	0	2	3

$$w_{u,v} = \begin{cases} \frac{2 \times \text{sim}(u,v) \times T_{u,v}}{\text{sim}(u,v) + T_{u,v}} & \text{if } \text{sim}(u,v) + T_{u,v} \neq 0 \\ T_{u,v} & \text{and } \text{sim}(u,v) \times T_{u,v} \neq 0 \\ \text{sim}(u,v) & \text{else if } \text{sim}(u,v) = 0 \text{ and } T_{u,v} \neq 0 \\ 0 & \text{else if } \text{sim}(u,v) \neq 0 \text{ and } T_{u,v} = 0 \\ 0 & \text{else} \end{cases} \quad (۵)$$

در واقع رابطه (۵) ترکیب روابط اعتماد و شباهت بین دو کاربر u و v است. در رابطه (۵) اگر دو معیار شباهت و اعتماد میان دو کاربر u و v غیر صفر باشد، وزن میان آن‌ها از تناسب دو برابر حاصل ضرب اعتماد در شباهت بر روی مجموع اعتماد و شباهت به دست می‌آید، اما اگر شباهت میان دو کاربر صفر باشد، وزن آن‌ها برابر با اعتماد میان دو کاربر و یا اگر اعتماد میان دو کاربر صفر باشد، وزن میان آن‌ها برابر با شباهت میان آن‌هاست. درواقع در این مرحله، مجموعه کاربران به یک گراف $G = (V, E, W)$ نگاشت می‌شوند، که V مجموعه کاربران را به عنوان رئوس گراف نشان می‌دهد، E و W به ترتیب مجموعه یال‌ها و وزن‌های شباهت بین هر جفت کاربر می‌باشند. گراف موردنظر می‌تواند بر اساس اطلاعات کاربران موجود در سیستم ساخته شود. مقدار وزن نهایی بین هر جفت کاربر با استفاده از رابطه (۵) قابل محاسبه است.

تعیین تعداد بهینه خوشه‌ها

یکی از مشکلات جدی سیستم‌های پیشنهاددهنده، شروع سرد است. مشکل شروع سرد برای کاربر جدید زمانی رخ می‌دهد که به علت نبود یا کمبود رتبه‌های اولیه امکان ارائه پیشنهادها مناسب وجود ندارد. یکی از راه‌های مقابله با این مشکل استفاده از تکنیک‌های خوشه‌بندی است که عموماً برای حل مشکل شروع سرد بکار گرفته می‌شوند. در این روش می‌توان کاربران را خوشه‌بندی کرد به گونه‌ای که کاربران با بیشترین شباهت در یک خوشه قرار گیرند.

در روش پیشنهادی بعد از محاسبه شباهت میان هر جفت کاربر، به خوشه‌بندی کاربران پرداخته می‌شود که شامل دو بخش تعیین تعداد بهینه خوشه‌ها و سپس خوشه‌بندی کاربران است. در ابتدا به دنبال تعیین گروه‌هایی از کاربران هستیم که از نظر علاقه‌مندی بیشترین شباهت را به یکدیگر دارند. روش خوشه‌بندی استفاده‌شده در این مقاله C-Means فازی است. روش خوشه‌بندی C-Means فازی برخلاف روش سنتی هرکدام از آیتم‌ها را با درجه‌های عضویت مختلف در خوشه‌های مختلف قرار می‌دهد و مجموع درجه عضویت‌های هر آیتم در C خوشه متفاوت برابر با ۱ است.

در روش سنتی خوشه‌بندی، هر آیتم تنها عضو یک مجموعه هست و می‌توان با قطعیت گفت که آیتم A عضو خوشه C هست یا خیر و اگر عضو این خوشه باشد دیگر عضو هیچ خوشه دیگری نیست. یعنی عضویت به صورت مطلق با ۰ یا ۱ به معنای عضو بودن در یک خوشه یا نبودن تعریف می‌شود. استفاده از روش سنتی باعث بروز خطاهایی می‌شود زیرا در دنیای واقعی معمولاً نمی‌توان مرز مشخصی بین دسته‌بندی‌های مختلف قرارداد. در روش خوشه‌بندی فازی C-Means

ماتریس اعتماد معمولاً بسیار خلوت است، لذا می‌توان با استفاده از روابط اعتمادی که در ادامه مطرح می‌شود میزان دقیق‌تر اعتماد هر کاربر به کاربر دیگر را با توجه به این مورد که اعتماد دارای خاصیت تعدی بین کاربران است محاسبه نمود.

محاسبه شباهت بین کاربران

در این قسمت، شباهت و اعتماد میان هر جفت کاربر بر اساس دو ورودی سیستم که ماتریس‌های رتبه و اعتماد می‌باشند محاسبه می‌شوند. شبکه اعتماد کاربر هدف می‌تواند بر اساس ترکیب روابط اعتماد و معیار ضریب همبستگی پیرسون به عنوان مقادیر نهایی شباهت، ساخته شود. در ادامه ابتدا اعتماد میان هر جفت کاربر محاسبه شده و پس از محاسبه شباهت میان آن‌ها، هر دو نتیجه بر اساس رابطه (۵) ترکیب شده و ماتریس کاربر-کاربر بر اساس مقدار به دست آمده تشکیل می‌شود.

از رابطه (۲) برای محاسبه اعتماد بین دو کاربر u و v استفاده می‌شود:

$$T_{u,v} = \frac{d_{\max} - d_{u,v} + 1}{d_{\max}} \quad (۲)$$

در رابطه (۲)، $T_{u,v}$ میزان اعتماد بین کاربر u و کاربر v ، $d_{\max}$ ماکزیم طول انتشار اعتماد بین کاربران و $d_{u,v}$ طول انتشار اعتماد بین کاربر u و کاربر v را نشان می‌دهد [۳۶]. مقدار $d_{\max}$ به صورت تقریبی، برابر با میانگین طول مسیر در شبکه اعتماد است که بر اساس رابطه (۳) محاسبه می‌شود [۴۹]. طول انتشار اعتماد برابر تعداد یال‌هایی است که در کوتاه‌ترین مسیر انتشار اعتماد، از کاربر اعتماد کننده به کاربر مورد اعتماد وجود دارد:

$$d_{\max} = [L^R] = \left\lceil \frac{\ln(n)}{\ln(k)} \right\rceil \quad (۳)$$

L^R میانگین طول مسیر در شبکه اعتماد، n اندازه شبکه اعتماد و k میانگین درجه در شبکه اعتماد را نشان می‌دهد. بر اساس رابطه (۲) هرچه طول انتشار بین دو کاربر کمتر باشد، میزان اعتماد آن‌ها بیشتر است. در ضمن، اگر طول انتشار اعتماد بین دو کاربر بیش از حداکثر طول انتشار در نظر گرفته شده باشد، اعتمادی میان آن دو کاربر وجود ندارد.

از طرفی، برای محاسبه شباهت میان کاربران از ماتریس رتبه و ضریب همبستگی پیرسون استفاده شده است که از رابطه (۴) به دست می‌آید:

$$\text{sim}(u, v) = \frac{\sum_{i \in A_{u,v}} (r_i(u) - \bar{r}(u))(r_i(v) - \bar{r}(v))}{\sqrt{\sum_{i \in A_{u,v}} (r_i(u) - \bar{r}(u))^2} \sqrt{\sum_{i \in A_{u,v}} (r_i(v) - \bar{r}(v))^2}} \quad (۴)$$

در رابطه (۴)، مقدار $r_i(u)$ رتبه داده‌شده به آیتم i توسط کاربر u ، $\bar{r}(u)$ میانگین رتبه‌های داده‌شده توسط کاربر u ، مجموعه $A_{u,v}$ مجموعه آیتم‌هایی است که توسط هر دو کاربر u و v رتبه داده شده‌اند و $\text{sim}(u, v)$ میزان شباهت میان کاربر u و v را نشان می‌دهد. سرانجام، وزن نهایی میزان شباهت و اعتماد بین کاربر u و کاربر v که با $w_{u,v}$ نشان داده می‌شود، بر اساس رابطه (۵) به صورت زیر محاسبه می‌گردد [۴۶]:

بعد از مشخص کردن تعداد خوشه‌های بهینه بر اساس X-Means، در مرحله بعد، خوشه‌بندی کاربران توسط خوشه‌بندی C-Means فازی انجام خواهد شد.

جدول ۴: تنظیمات الگوریتم X-Means
Table 4: Configurations of the X-Means Algorithm

پارامترهای پیکربندی Configuration parameters	مقدار Value
حد پایین	۲
Low bound (K_{min})	
حد بالا	۶۰
High bound (K_{max})	
بیشترین تعداد تکرار الگوریتم K-Means Maximum number of K-Means algorithm iterations	۱۰
بیشترین تعداد کل تکرارها Maximum total number of iterations	۱۰۰

خوشه‌بندی فازی کاربران و کشف قوانین انجمنی وزن دار

در خوشه‌بندی C-Means فازی، اگر مجموعه کاربران را $U = [u_1, u_2, \dots, u_n]$ و مجموعه خوشه‌ها را $C = [c_1, c_2, \dots, c_c]$ در نظر بگیریم، هر عضو مجموعه U می‌تواند با یک مقدار بین ۰ تا ۱ به عنوان درجه عضویت در بیش از یک خوشه قرار بگیرد. بنابراین، با توجه به این که هر کاربر در بیش از یک خوشه قرار می‌گیرد، یکی از چالش‌های سیستم‌های پیشنهاددهنده که خلوت بودن داده‌ها است تا حدود زیادی حل می‌شود و همسایگان بیشتری برای کاربر وجود خواهد داشت. از طرف دیگر با توجه به خوشه‌بندی و این که محاسبات تخمین رتبه تنها بر اساس همسایه‌های هر کاربر انجام می‌شود، روش پیشنهادی در حجم انبوه اطلاعات عملکرد مناسب‌تری خواهد داشت. شبه کد مراحل انجام کار C-Means فازی مطابق الگوریتم ۲ است:

الگوریتم ۲: شبه کد الگوریتم C-Means فازی
Algorithm 2: Pseudocode of the fuzzy C-Means algorithm

The pseudo-code of the Fuzzy C-Means (FCM)
Input: X and C.
Output: Final Fuzzy C-Means clusters.
1: Select an initial fuzzy pseudo-partition, i.e., assign values to all w_{ij}
2: Repeat
3: Compute the centroid of each cluster using the fuzzy partition.
4: Update the fuzzy partition, i.e., the w_{ij} .
5: Until the centroids don't change.

همان‌طور که در شبه کد بالا مشاهده می‌شود، اگر U مجموعه n کاربر و C مجموعه خوشه i و j به ترتیب شاخص‌های هر یک از کاربران و خوشه‌ها در نظر گرفته شود، ابتدا میزان تعلق هر کاربر به خوشه‌ها که بر اساس $w_{i,j}$ محاسبه می‌شود، انجام شده و سپس مرکز هر خوشه مشخص می‌شود. سپس به طور مجدد خوشه‌بندی کاربران با توجه به

می‌توان هر آیتم را با توجه به مشخصات آن در خوشه‌های مختلفی با درجه‌های عضویت مختلف قرار داد. این الگوریتم دارای نقاط قوت همچون همگرا بودن و بدون نظارت بودن و نقاط ضعف از جمله زمان محاسبات زیاد، حساس به حدس‌های اولیه و امکان توقف در نقاط بهینه محلی است.

یکی از مشکلات الگوریتم‌های خوشه‌بندی، حدس اولیه در تعیین تعداد خوشه‌ها پیش از اجرای الگوریتم است. بنابراین قبل از آن که به خوشه‌بندی فازی C-Means در روش پیشنهادی بپردازیم باید تعداد بهینه خوشه‌ها را تعیین کنیم. به این منظور و برای تعیین تعداد بهینه خوشه‌ها از الگوریتم X-Means استفاده شده است [۴۴]. این الگوریتم توسعه‌یافته الگوریتم K-Means است و همانطور که در بخش مفاهیم اولیه توضیح داده شده، سعی دارد به‌طور خودکار تعداد خوشه‌ها را بر اساس امتیازهای معیار اطلاعات بیزین مطابق با رابطه (۱) تعیین کند. این الگوریتم در ابتدا با مقداری برابر K برای حد پایین تعداد مراکز شروع می‌کند و مراکز را اضافه کرده تا به حد بالا برسد. در طول این فرآیند مراکز برای دستیابی به بهترین امتیاز تغییر می‌کنند. شبه کد X-Means در الگوریتم ۱ ارائه شده است. همچنین تنظیم پارامترهای این الگوریتم را می‌توان در جدول ۴ مشاهده کرد.

الگوریتم ۱: شبه کد الگوریتم X-Means
Algorithm1: The pseudocode of the X-Means algorithm

The Pseudo-Code Of X-Means Clustering
Input: a dataset D, upper bounds (K_{max}) and lower bound (K_{min}) for possible K
Output: an Integer k as the number of clusters
1. Initialize $K = K_{min}$
2. Clusters $\leftarrow$ Run K-means algorithm on D
3. While $K < K_{max}$
 {New Clustering $\leftarrow$ {
4. For $\forall C_i \in$ Clusters
 {
5. Children Structure $\leftarrow$ Split C_i into 2 children clusters
6. If $BIC(\text{Children Structure}) > BIC(C_i)$
7. New Clustering $\leftarrow$ Children Structure
8. Else
9. NewClustering $\leftarrow C_i$
10. }
11. Clusters $\leftarrow$ NewClustering
12. If $BIC(\text{Clusters}) > \text{Best BIC obtained from previous steps}$
13. Best Structures $\leftarrow$ Clusters
14. }
15. Return |Best Structures|

بر اساس [۴۴] که در آن جزییات نحوه اجرای الگوریتم X-Means شرح داده شده است، مقدار K_{min} و K_{max} که به ترتیب کمترین و بیشترین تعداد مرکز خوشه‌ها هستند و توسط کاربر مقداردهی اولیه می‌شوند به ترتیب برابر با اعداد ۲ و ۶۰ تعیین شده و همچنین بیشترین تعداد تکرار الگوریتم K-Means برابر با ۱۰ و بیشترین دفعات کل اجرا برابر با ۱۰۰ است.

موردنظر و وزن متناظر آن‌ها تشکیل شده است. ساختار زوج مرتب‌های بردار L_k برای کاربر k ام به صورت

$$L_k = ((e_1, w(e_{1k})), ((e_2, w(e_{2k}))), \dots, ((e_m, w(e_{mk}))))$$

می‌شوند که $w(e_{mk})$ نشان‌دهنده وزن دوره آموزشی e_m در لیست کاربر u_k است. در واقع وزن هر دوره آموزشی بارتبه‌ای که کاربر به دوره آموزشی e_1 داده است مقداردهی می‌شود. در این بردارها، الگوهای دسترسی کاربران به دوره‌های آموزشی شناسایی شده و برای آن‌ها وزن در نظر گرفته می‌شود. اگر الگوها را X در نظر بگیریم، وزن این الگوها که بر اساس علاقه‌مندی کاربر به آیت‌های تشکیل دهنده آن‌ها محاسبه می‌شود، با عبارت $w(X, L)$ نشان داده می‌شود. ساده‌ترین راه برای بدست آوردن وزن، در نظر گرفتن مینیمم وزن آیت‌های تشکیل دهنده الگو می‌باشد و در رابطه (۹) نشان داده شده است:

$$w(X, L) = \begin{cases} \min(w(e_1, e_2, \dots, e_s)) & X \subseteq L \\ 0 & X \not\subseteq L \end{cases} \quad (9)$$

در رابطه (۹) X نشان دهنده یک الگو، L لیست زوج مرتب‌های دوره‌های آموزشی گذرانده شده کاربر و وزن آن‌هاست. با اختصاص وزن به الگوها، می‌توان برای هر یک از لیست‌های دوره‌های آموزشی متعلق به هر کاربر نیز وزن تعریف کرد. به دلیل آن که لیست دوره‌های آموزشی هر کاربر، با درجه‌های متفاوتی به خوشه‌های فازی مختلف، متعلق است، بنابراین وزن و ارزش هر کاربر در خوشه‌های متفاوت، به درجه تعلق لیست آن کاربر به خوشه‌ها بستگی دارد و وزن نهایی کاربر u_k در خوشه C معادل با وزن نهایی L_k (لیست دوره‌های آموزشی کاربر k ام) بوده و مطابق با رابطه (۱۰) محاسبه می‌شود:

$$w(L_k) = w(u_k) = S_{kc} * \frac{\sum_{i=1}^{|L_k|} w(e_{ik})}{|L_k|} \quad (10)$$

در رابطه (۱۰) S_{kc} درجه تعلق لیست مربوط به کاربر u_k به خوشه C و $w(e_{ik})$ وزن اختصاص داده شده به دوره آموزشی e_i در لیست کاربر u_k است. ضریب پشتیبانی قوانین انجمنی وزن‌دار الگوها در میان لیست‌های همه کاربران، مطابق رابطه (۱۱) محاسبه می‌گردد. برای به دست آوردن ضریب پشتیبانی قوانین انجمنی وزن‌دار الگو X ، باید نسبت حاصل ضرب وزن لیست هر کاربر و وزن الگو X در لیست مربوط آن کاربر را به حاصل ضرب وزن لیست‌های همه کاربران در میانگین وزن همه دوره‌های آموزشی به دست آوریم:

$$wsp(X) = \frac{\sum_{u_i \in U} w(L_i) * w(X, L_i)}{\bar{w} * \sum_{j=1}^{|U|} w(L_j)} \quad (11)$$

در رابطه (۱۱) $\bar{w}$ میانگین وزن همه دوره‌های آموزشی در کل لیست‌های کاربران و U مجموعه کاربران است. ضریب اطمینان وزن‌دار برای قوانین انجمنی وزن‌دار بر اساس رابطه (۱۲) تعریف می‌شود:

$$wconf(X \Rightarrow Y) = \frac{wsp(X \cup Y)}{wsp(X)} \quad (12)$$

در این مدل از قوانین انجمنی، علاوه بر پارامترهای ضریب پشتیبانی و ضریب اطمینان وزن‌دار، وزن متناظر هر دوره آموزشی گذرانده شده

مراکز جدید انجام می‌شود. این کار تا زمانی که تغییرات مراکز خوشه‌ها از یک حد آستانه مشخص کوچکتر شود، تکرار شده و در نهایت خوشه‌های نهایی به دست می‌آیند.

پس از دریافت تعداد خوشه‌ها از الگوریتم X-Means، خوشه‌بندی C-Means فازی انجام می‌شود. کاربران سیستم بر اساس سلیقه و علاقه‌مندی‌هایشان خوشه‌بندی شده و یک ماتریس عضویت حاصل می‌شود که نشان‌دهنده میزان عضویت هر کاربر در هر خوشه است. این ماتریس با C سطر و n ستون حاوی مؤلفه‌هایی با مقداری بین [۰ تا ۱] است. تابع هدفی که در الگوریتم C-Means فازی دنبال می‌شود مطابق رابطه (۶) است:

$$J = \sum_{i=1}^n \sum_{j=1}^c S_{ij}^m \|x_i - v_j\|^2 \quad (6)$$

$$v_j = \frac{\sum_{i=1}^n S_{ij}^m x_i}{\sum_{i=1}^n S_{ij}^m} \quad (7)$$

$$S_{ij} = \frac{1}{\sum_{k=1}^c \left(\frac{d_{ij}}{d_{ik}} \right)^{\frac{2}{(m-1)}}} = \frac{1}{\sum_{k=1}^c \left(\frac{\|x_i - v_j\|}{\|x_i - v_k\|} \right)^{\frac{2}{(m-1)}}} \quad (8)$$

در رابطه (۶)، m یک عدد حقیقی بزرگ‌تر از ۱ است که در اکثر موارد برای m عدد ۲ انتخاب می‌شود. اگر m را برابر ۱ قرار دهیم تابع هدف خوشه‌بندی C-Means (کلاسیک) غیرفازی به دست می‌آید. x_i نمونه i ام است و v_j مرکز خوشه j ام است که با استفاده از رابطه (۷) محاسبه می‌شود. S_{ij} میزان تعلق نمونه i به خوشه j ام است که بر اساس رابطه (۸) محاسبه می‌شود. علامت $\|*\|$ نشان‌دهنده میزان فاصله نمونه از مرکز خوشه است که برای آن می‌توان از هر تابع فاصله‌ای استفاده کرد و در این مقاله از تابع فاصله اقلیدسی استفاده شده است. بر اساس S_{ij} ها ماتریس S تعریف می‌شود که دارای C سطر و n ستون است و مؤلفه‌های آن مقادیر بین ۰ تا ۱ دارند. مجموع مؤلفه‌های هریک از ستون‌ها باید برابر ۱ باشد.

بعد از خوشه‌بندی باید از طریق کاوش قوانین انجمنی، مجدداً کاربران در خوشه‌ها وزن‌دهی شوند تا به این صورت درجه اهمیت عضویت هر کاربر در خوشه‌های مختلف به دست بیاید. قوانین انجمنی ارتباطات میان کاربران را بر مبنای الگوهای انتخاب دوره‌های آموزشی کاربران و میزان علاقه‌مندی آن‌ها به دوره‌های آموزشی، بدون در نظر گرفتن ترتیب آن‌ها نشان می‌دهد. بیشتر رویکردهای کشف قوانین انجمنی بر مبنای الگوریتم APRIORI می‌باشند [۵۰]. در این مرحله، باید قوانین انجمنی وزن‌دار هر خوشه استخراج گردد. هدف الگوریتم پیشنهادی از کاوش قوانین انجمنی با اختصاص وزن به هر کاربر در خوشه‌ها، نشان دادن میزان اهمیت آن کاربر در آن خوشه خاص است.

U مجموعه کاربران سیستم $U = \{u_1, u_2, \dots, u_n\}$ و E مجموعه دوره‌های آموزشی ارائه شده $E = \{e_1, e_2, \dots, e_m\}$ در سیستم در نظر گرفته می‌شود. سپس برای هر کاربر یک بردار m بعدی به نام L تعریف می‌کنیم که از زوج‌های مرتب حاوی دوره‌های آموزشی گذرانده شده توسط کاربر

در رابطه (۱۵) $w(e_{ik})$ وزن اختصاص داده شده به دوره آموزشی e_i توسط کاربر k ام و $w(e_i)$ وزن دوره آموزشی e_i در سرایند قانون انجمنی R_L است. به دلیل آنکه هدف از شخصی سازی کردن سیستم پیشنهاددهنده، محاسبه یک مجموعه پیشنهادی با بیشترین مطابقت با علاقه مندی های یک کاربر است، باید قابلیت پیشنهاد دادن هر یک از دوره های آموزشی e_i که کاربر نگذرانده است، بر اساس قوانین انجمنی محاسبه شود. بنابراین بر اساس رابطه (۱۶)، با بهره گیری از درجه تطابق و ضریب اطمینان وزن دار، برای تمامی دوره های آموزشی گذرانده نشده توسط کاربر، امتیاز پیشنهاد محاسبه می شود. هر چه این امتیاز بالاتر باشد، احتمال آن که دوره آموزشی مربوطه مورد علاقه کاربر باشد، بیشتر است:

$$\text{Recommendation Score } (w_k, X \Rightarrow e_i) = \text{Match Score}(w_k, X) * wconf(X \Rightarrow e_i) \quad (16)$$

در نهایت، بر اساس قوانین انجمنی وزن دار، لیستی از دوره های آموزشی با بالاترین امتیاز برای پیشنهاددهی به کاربر کاندید می شوند و در لیستی به نام candidate ذخیره می شوند.

پیشنهاد دوره های آموزشی به کاربر

در بخش نهایی بر اساس امتیازهای محاسبه شده در مرحله قبل، دوره های آموزشی به صورت نزولی مرتب شده و N دوره آموزشی با بیشترین رتبه به عنوان آیتم های مورد علاقه کاربر هدف به وی پیشنهاد داده می شوند.

شبه کد نهایی روش پیشنهادی مبتنی بر خوشه بندی فازی و روابط اعتماد در الگوریتم ۳ ارائه شده است. خروجی الگوریتم پیشنهادی مقدار top_N نشان دهنده تعداد آیتم های آموزشی لیست پیشنهاد است. در خط اول شبه کد، مجموعه داده ورودی به دو مجموعه آموزش و تست تقسیم می شود. بدین منظور، ۸۰ درصد از رتبه های هر کاربر به صورت تصادفی به عنوان داده ی آموزش و مابقی رتبه ها به عنوان داده تست در نظر گرفته می شوند. داده آموزش به منظور ساخت مدل سیستم پیشنهادی مورد استفاده قرار می گیرد، در حالی که از داده تست برای ارزیابی روش پیشنهادی استفاده می شود. در خط دوم شبه کد، شباهت بین کاربران بر اساس داده آموزش و روابط اعتماد و شباهت محاسبه می شود. در خطوط (۳ تا ۵) به ترتیب تعیین تعداد خوشه ها با استفاده از الگوریتم X-Means، خوشه بندی کاربران با استفاده از خوشه بندی

C-Means فازی و تکرار مراحل خوشه بندی با یافتن مراکز خوشه جدید تا زمانی که دیگر مراکز خوشه ها تغییر نکنند، انجام می شود. پس از تعیین خوشه های نهایی، رتبه های مورد نظر برای کاربر هدف بر اساس قوانین انجمنی وزن دار پیش بینی شده و تعداد top_N آیتم مورد علاقه کاربر به وی پیشنهاد می شود (خطوط ۶-۱۱).

توسط کاربران نیز نشان داده می شود. قوانین انجمنی وزن دار به صورت رابطه (۱۳) تعریف می شوند:

$$R = \langle (e_1, e_2, \dots, e_k), (q_{k+1}, q_{k+2}, \dots, q_{k+p}), (w_1, w_2, \dots, w_{k+p}), \delta, \alpha \rangle \in R \quad (13)$$

در رابطه (۱۳) به ترتیب $(q_{k+1}, q_{k+2}, \dots, q_{k+p})$ و $(e_1, e_2, \dots, e_k)$ نشان دهنده سرایند و بدنه قانون انجمنی، δ بیان گر ضریب پشتیبانی وزن دار، α ضریب اطمینان وزن دار و $(w_1, w_2, \dots, w_{k+p})$ وزن متناظر با هر یک از دوره های آموزشی است. در واقع خروجی این مرحله، استخراج قوانین فازی وزن دار برای هر یک از خوشه های فازی و ایجاد یک ماتریس عضویت برای تعیین درجه عضویت هر کاربر در هر خوشه می باشد.

انتخاب دوره های آموزشی کاندید بر اساس قوانین انجمنی و پیش بینی رتبه

با توجه به روش خوشه بندی فازی، کاربر هدف دارای درجه عضویت مختلف در خوشه ها است و طیف فازی مناسب برای او انتخاب شده است. در این مرحله برای تعیین دوره های آموزشی انتخابی از هر خوشه باید یک تناسب میان تعداد دوره های آموزشی از هر خوشه و درجه تعلق کاربر به آن خوشه برقرار شود. مثلاً اگر بخواهیم ۱۰ دوره آموزشی به کاربر پیشنهاد دهیم و درجه تعلق کاربر به یک خوشه برابر با ۰/۷ و به خوشه دیگر ۰/۳ باشد، تعداد دوره های آموزشی انتخابی از خوشه اول هفت عدد و از خوشه دوم سه عدد خواهد بود.

سپس لیست دوره های آموزشی گذرانده شده توسط کاربران با سرایند قوانین انجمنی وزن دار مربوط به هر خوشه فازی مقایسه شده و قوانینی که دوره های آموزشی بخش سرایند آن ها در لیست دوره های گذرانده شده توسط کاربر وجود دارد، استخراج می شوند. سپس مطابقت وزن دوره های آموزشی گذرانده شده توسط کاربر k ام به نام w_k که عبارت است از $(w(e_{1k})), (w(e_{2k})), \dots, (w(e_{mk}))$ با وزن آیتم های سرایند قوانین انجمنی وزن دار مربوط به هر خوشه فازی $(w(e_m)), (w(e_2)), \dots, (w(e_1)) = R_L$ مورد بررسی قرار می گیرد. رابطه (۱۴) درجه مطابقت لیست دوره های آموزشی گذرانده شده کاربر با هر یک از قوانین انجمنی وزن دار استخراج شده، را محاسبه می کند:

$$\text{MatchScore}(w_k, R_L) = 1 - \frac{1}{4} \sqrt{\frac{\text{Dissimilarity}(w_k, R_L)}{\sum_{i: w(e_i) > 0} 1}} \quad (14)$$

در رابطه (۱۴) $w(e_i)$ وزن دوره آموزشی e_i در سرایند قانون انجمنی R_L است و مقدار Dissimilarity در رابطه (۱۴) که نشان دهنده عدم تطابق سرایند قانون انجمنی وزن دار و دوره های آموزشی گذرانده شده کاربر است بر اساس رابطه (۱۵) محاسبه می شود.

$$\text{Dissimilarity}(w_k, R_L) = \sum_{i: w(e_i) > 0} \frac{2 * (w(e_{ik}) - w(e_i))}{(w(e_{ik}) + w(e_i))} \quad (15)$$

برچسب نمایش داده خواهد شد. برای پیاده‌سازی روش پیشنهادی، از جهت نرم‌افزاری از زبان برنامه‌نویسی Matlab و از نظر سخت‌افزاری نیز، برنامه‌های نوشته‌شده بر روی یک سیستم با هسته پردازشی Corei5 و ۴ گیگابایت حافظه داخلی (RAM) اجرا شده‌اند. تعداد دفعات اجرای کد ۱۰۰ بار است و داده‌ها یک بار شافل شده‌اند. مجموعه داده‌ها به‌طور تصادفی به ۸۰٪ آموزش و ۲۰٪ تست تقسیم شده‌اند و از مجموعه آموزش برای تولید مدل پیشنهادی استفاده گردیده است.

مجموعه داده

در این قسمت به معرفی کامل مجموعه داده Moodle می‌پردازیم و فایل‌های موجود در این مجموعه داده را شرح می‌دهیم. اطلاعات موجود در این مجموعه داده مربوط به برگزاری یک دوره‌ی آموزشی به‌صورت MOOC (Massive Open Online Course) و به طول ۴ هفته به نام "Teaching with Moodle" است که سالانه ۲ بار به‌صورت کاملاً برخط و رایگان برگزار می‌شود که در آن شرکت‌کنندگان می‌توانند به‌طور کامل با روش‌های یادگیری الکترونیکی از طریق بستر Moodle آشنا شوند و بعد از اتمام این دوره اقدام به راه‌اندازی سیستم خود و آموزش از طریق Moodle نمایند. مجموعه داده مورد استفاده مربوط به برگزاری دوره در ماه اوت سال ۲۰۱۶ است که در آن ۶۱۱۹ کاربر از طریق به‌صورت برخط اقدام به شرکت در این دوره کرده‌اند و بعد از اتمام دوره با توجه به رضایت کاربران داده‌های مربوط به ۲۱۶۸ کاربر به‌صورت ناشناس برای تحقیقات به‌صورت عمومی منتشر شده است. این مجموعه داده که شامل ۶ فایل CSV به صورت جداول سطری-ستونی است را می‌توان به صورت برخط دریافت کرد. در ادامه هر یک از این فایل‌ها توضیح داده می‌شوند. در جدول ۵ می‌توان فایل‌های این مجموعه داده را به همراه توضیحات و برخی از ستون‌ها مشاهده کرد.

الگوریتم ۳: شبه‌کد روش پیشنهادی مبتنی بر خوشه‌بندی فازی و روابط اعتماد

Algorithm 3: Pseudocode of the proposed method based on fuzzy clustering and trust relations

The pseudo-code of the proposed methods

Input: user set, coarse set, ratings, grades

Output: Top-N recommendation list.

- 1: Split dataset into train set Tr and test set Te ;
- 2: Calculate user similarities using the train set Tr with pearson and trust equations;
- 3: Determine the number of clusters with x-means algorithm
- 4: Apply the Fuzzy C-means algorithm
- 5: Update cluster centers, until cluster centers do not change
- 6: For all users $u \in Te$ do
- 7: Calculation of the degree of the membership of the user to each cluster
- 8: Finding weighted association rules
- 9: Calculation of candidate courses recommendation score based on Eq. (16)
- 10: End.
- 11: Recommend top_N of items as the recommendation list to the user u ;

نتایج و بحث

در این بخش، نتایج آزمایش‌های انجام‌شده برای ارزیابی روش پیشنهادی بر روی مجموعه داده Moodle ارائه می‌شود. برای این منظور، ابتدا، مجموعه داده Moodle به‌طور کامل تشریح می‌شود، در ادامه معیارهای ارزیابی معرفی می‌شوند و سپس آزمایش‌های انجام‌شده و نتایج آن‌ها گزارش می‌شوند. همچنین برای مقایسه نتایج حاصل از روش پیشنهادی با سایر روش‌ها و کارهای پیشین، با پیاده‌سازی روش‌های دیگر با توجه به مجموعه داده مورد استفاده در این مقاله، بررسی و مقایسه دقیق بین روش‌ها صورت می‌پذیرد.

روش پیشنهادی به اختصار (Trust-Based Recommender System for TBEL E-Learning) نام‌گذاری می‌شود و در ادامه نتایج حاصل با این

جدول ۵: مشخصات فایل‌های موجود در مجموعه داده Moodle

Table 5: Characteristics of files related to the Moodle dataset

نام فایل File name	ستون‌ها Columns (fields)	تعداد سطرها Row numbers	توضیحات Description
mdl_badge_issued	UserName UniqueHash DateIssued Visibility	2354	این فایل حاوی همه مدال‌هایی است که کاربران در دوره‌های آموزشی مختلف کسب کرده‌اند This file contains all the badges that users have been able to obtain in different educational courses.
mdl_course_modules	CourseName CourseID Module Instance Section Completion Score	75	این فایل حاوی آیتم‌هایی است که مربوط به دوره‌های آموزشی هستند و کاربران در آن‌ها شرکت کرده‌اند This file contains items that are related to the course and users can participate in them
mdl_course_modules_completion	UserName UserID CourseModuleID CompletionState Time	41204	این فایل نشان می‌دهد که هر کاربر چه دوره‌های آموزشی را گذرانده است This file shows the educational courses that each user has taken.

mdl_grade_grades_history	UserName UserID ItemID Time RawGrade RawGradeMax RawGradeMin RawScaleID FinalGrade	87738	این جدول شامل تاریخچه عملکردی کامل کاربران در آیتم‌های دوره‌های آموزشی است مانند امتیاز کسب‌شده توسط هر کاربر و همچنین نظر کاربران در مورد آیتم پشت سر گذاشته شده This file contains the complete performance history of the users in the items of educational courses. For example, the grades of each user or the users' opinions about the passed items.
mdl_logstore_standard_log	UserName EventName Component ItemID Action Time	346584	این جدول شامل تمامی لیست‌های ثبت‌شده از تمامی رخداد-هایی است که رفتار کاربران در سیستم را نشان می‌دهد This file contains the recorded lists of all the events that present the behaviour of the users
mdl_user	UserName UserID FirstAccess LastAccess Country City Language TimeZone BadgesCount	2170	این فایل حاوی اطلاعات پروفایل کاربر است This file contains user profile information

که کاربران توانسته‌اند دریافت کنند. بر اساس داده‌های این فایل، تعداد مدال‌های هر کاربر مشخص شده و جدول Badges ایجاد می‌شود که شامل فیلدهای نشان داده شده در جدول ۷ است:

جدول ۷: مشخصات فیلدهای جدول Badges
Table 7: Description of Badges table's fields

فیلد Field	توضیحات Description
UserName	نام کاربری هر کاربر Username for each user
UniqueHash	شناسه هر افتخار Identifier for each badge
DateIssued	زمان دریافت مدال توسط کاربر Time of receiving the badge by the user

معیارهای ارزیابی

در این بخش به معرفی معیارهای ارزیابی سیستم پیشنهادی در این مقاله می‌پردازیم. معمولاً برای آزمون سیستم‌های پیشنهاددهنده از روش اعتبارسنجی یک‌طرفه استفاده می‌شود. در این روش، سیستم بدون در نظر گرفتن امتیازی که کاربر جاری به آیتم هدف داده است، این امتیاز را با استفاده از همسایگان این کاربر و امتیازهایی که او به سایر آیتم‌ها داده است، پیش‌بینی می‌کند. سپس با استفاده از معیارهای ارزیابی، امتیاز پیش‌بینی‌شده را با امتیاز واقعی که کاربر جاری به آیتم هدف داده است، مقایسه می‌کند. در ادامه، معیارهایی که در این مقاله برای ارزیابی و مقایسه TBEL استفاده شده است، توضیح داده می‌شوند.

معیارهای کیفیت پیش‌بینی

پیش‌بینی، یک مقدار عددی است که سیستم پیشنهاددهنده به عنوان رتبه کاربر u به آیتم i برمی‌گرداند. معمولاً سیستم‌های پیشنهاددهنده،

با توجه به اطلاعات بالا در مورد تعداد سطرهای هر فایل اگر تمامی ۲۱۷۰ کاربر به‌تمامی ۷۵ دوره آموزشی موجود رأی داده باشند باید تعداد ۱۶۲۷۵۰ رکورد در فایل mdl_grades_history ثبت شده باشد، اما به دلیل اینکه تمامی کاربران در تمامی دوره‌ها شرکت نکرده‌اند و به‌تمامی آن‌ها رأی نداده‌اند، لذا ۸۷۷۳۸ رکورد در این فایل وجود دارد. همان‌طور که در بخش‌های ایجاد ماتریس رتبه و اعتماد شرح داده شد، دو ماتریس رتبه و اعتماد به عنوان ورودی‌های سیستم در نظر گرفته می‌شود که برای ساخت آن‌ها از فایل‌های mdl_user، mdl_grade_grades_history و mdl_badge_issued استفاده می‌شود. برای ایجاد ماتریس رتبه، تعدادی از فیلدهای فایل mdl_grade_grades_history استخراج می‌شود که دربرگیرنده اطلاعات مشارکت کاربران در دوره آموزشی است. فیلدهای جدول Grades در جدول ۶ نشان داده شده است:

جدول ۶: مشخصات فیلدهای جدول Grades
Table 6: Description of Grades table's fields

فیلد Field	توضیحات Explanations
UserName	نام کاربری هر کاربر Username for each user
UserID	شناسه کاربری هر کاربر User identifier for each user
ItemID	شناسه دوره آموزشی Course identifier
Rate (FinalGrade)	نمره هر کاربر برای دوره‌های آموزشی (بین ۱ تا ۵) Grade of each user for courses (between 1 and 5)

برای ایجاد ماتریس اعتماد، تعدادی از فیلدهای فایل mdl_badge_issued استخراج می‌شود که درواقع شامل افتخاراتی است

$|U'|$ ، تعداد کاربرانی است که برای آن‌ها حداقل یک رتبه پیش‌بینی شده و $|U|$ تعداد کل کاربران است.

ارزیابی روش پیشنهادی

اکنون، نتایج به دست آمده برای TBEL ارائه شده و آن‌ها را به منظور ارزیابی روش ارائه شده با روش‌های دیگر مورد بحث و بررسی قرار می‌دهیم. برای بررسی TBEL، مجموعه داده Moodle به‌طور تصادفی به دو قسمت ۸۰٪ آموزش و ۲۰٪ تست تقسیم می‌شود. به‌منظور سنجش میزان خطا و پوشش TBEL نسبت به سایر روش‌ها، از معیارهای MAE، RMSE، RC و UC استفاده شده است. نتایج حاصل از TBEL با ۴ روش TBHR[4]، NPR_eL[19]، LSEL[5] و CF مقایسه شده است. روش‌های ارایه شده در مقالات [۴] و [۱۹] بر اساس مجموعه داده Book-Crossing و روش پیشنهادی در مقاله [۵] نیز بر اساس مجموعه داده Moodle ارزیابی شده است. ما با هدف ارزیابی صحیح مقالات و امکان مقایسه روش پیشنهادی با روش‌های مورد مقایسه در شرایط یکسان، روش‌های TBHR، LSEL، NPR_eL و CF را شبیه‌سازی کرده و بر اساس مجموعه داده Moodle معرفی شده در بخش مجموعه داده، مورد ارزیابی قرار داده‌ایم. TBHR از روش جدیدی برای ترکیب اعتماد و خوشه‌بندی ساده K-Means در محیط‌های یادگیری الکترونیکی استفاده کرده است [۴]. از آنجا که در مقاله مقایسه شده نتایج بروی مجموعه داده متفاوتی انجام شده است، لذا روش TBHR با استفاده از مجموعه داده Moodle پیاده‌سازی شده است.

LSEL که در [۵] معرفی شده است یک سیستم پیشنهاددهنده یادگیری الکترونیکی است که پیشنهادها را شخصی‌سازی شده به کاربران ارائه می‌دهد. این روش برای ساخت پیشنهادها از سبک یادگیری مورد علاقه هر کاربر و بررسی مشخصات دوره‌های یادگیری ارائه شده استفاده می‌کند. در LSEL برای داشتن بهترین نتیجه، روش پیشنهادی با انواع الگوها برای یافتن شباهت مانند ضریب همبستگی پیرسون پیاده‌سازی شده است که در اینجا بهترین نتیجه کسب شده در LSEL برای معیارهای MAE و RMSE مورد مقایسه قرار گرفته است.

در [۱۹] نیز روشی به نام NPR_eL معرفی شده است که از ترکیب روش‌های پایه CF و محتوا محور برای ایجاد پیشنهاد دوره‌ها بر اساس سبک یادگیری کاربر استفاده می‌کند. از نتایج حاصل از این روش بر روی مجموعه داده Moodle برای مقایسه RC و UC استفاده شده است. علاوه بر این سه روش، نتایج به دست آمده با روش پایه‌ای CF نیز مقایسه شده است. لازم به ذکر است که CF به عنوان کلاسیک‌ترین روش مطرح شده در سیستم‌های پیشنهاددهنده تنها بر اساس معیار ضریب همبستگی پیرسون و ماتریس رتبه اقدام به تولید پیشنهاد می‌کند و این روش نیز بر روی مجموعه داده Moodle پیاده‌سازی شده است.

شکل ۲ نشان‌دهنده میزان MAE و RMSE روش پیشنهادی TBEL و سه روش دیگر است. نتایج نشان داده شده بیان‌گر این مطلب است که روش پیشنهادی دارای خطای کمتری به نسبت سه روش TBHR، LSEL و CF

بر اساس معیار دقت ارزیابی می‌شوند. این معیار اختلاف مقادیر پیش‌بینی شده را نسبت به مقدار واقعی آن‌ها محاسبه می‌کند. معیارهای کیفیت پیش‌بینی که در این مقاله برای ارزیابی روش‌های پیشنهادی مورد استفاده قرار گرفته‌اند، عبارت‌اند از:

میانگین خطای مطلق (MAE): این معیار، میانگین خطای مطلق رتبه‌های پیش‌بینی شده را به منظور اندازه‌گیری دقت سیستم پیشنهاد دهنده محاسبه می‌کند. برای این منظور، قدر مطلق تفاضل رتبه پیش‌بینی شده برای هر زوج کاربر-آیتم، با رتبه واقعی اندازه‌گیری می‌شود. میانگین خطای کل امتیازهای پیش‌بینی شده، به‌عنوان میانگین خطای مطلق در نظر گرفته می‌شود که با استفاده از رابطه (۱۷) به دست می‌آید [۱]:

$$MAE = \frac{\sum_{i=1}^N |r_i - p_i|}{N} \quad (17)$$

خطای جذر میانگین مربعات (RMSE): یک معیار پرکاربرد در کنار میانگین خطای مطلق، خطای جذر میانگین مربعات است و تأکید بیشتری بر روی خطاهای بزرگ‌تر دارد. این معیار با استفاده از رابطه (۱۸) محاسبه می‌گردد [۱]:

$$RMSE = \sqrt{\frac{1}{N} \sum_{i=1}^N (r_i - p_i)^2} \quad (18)$$

در روابط فوق، r_i رتبه واقعی کاربر و p_i رتبه پیش‌بینی شده و N تعداد کل پیش‌بینی‌های انجام‌شده توسط سیستم است.

معیارهای پوشش رتبه و کاربر

یکی دیگر از انواع معیارهای ارزیابی مهم که در سیستم‌های پیشنهاد دهنده مورد استفاده قرار می‌گیرند، معیارهای پوشش رتبه و کاربر می‌باشند. این معیارها درصد رتبه‌های پیش‌بینی شده و همچنین درصد کاربرانی که سیستم توانسته است برای آن‌ها رتبه‌ای پیش‌بینی کند را نشان می‌دهند. در ادامه دو معیار ارزیابی برای نرخ پوشش رتبه و کاربر که در این مقاله مورد استفاده قرار گرفته است، توضیح داده خواهند شد: نرخ پوشش رتبه‌ها (RC): این معیار درصد رتبه‌هایی که سیستم به درستی قادر به پیش‌بینی آن‌ها بوده است را محاسبه می‌کند. روش محاسبه این معیار بر اساس رابطه (۱۹) است [۱]:

$$RC = \frac{|I'|}{|I|} \quad (19)$$

$|I'|$ ، تعداد رتبه‌هایی است که سیستم آن‌ها را پیش‌بینی کرده است و $|I|$ ، تعداد کل رتبه‌های موجود در سیستم است. نرخ پوشش کاربر (UC): این معیار درصد کاربرانی را نشان می‌دهد که سیستم قادر است برای آن‌ها حداقل یک رتبه پیش‌بینی کند. روش محاسبه این معیار بر اساس رابطه (۲۰) است [۱]:

$$UC = \frac{|U'|}{|U|} \quad (20)$$

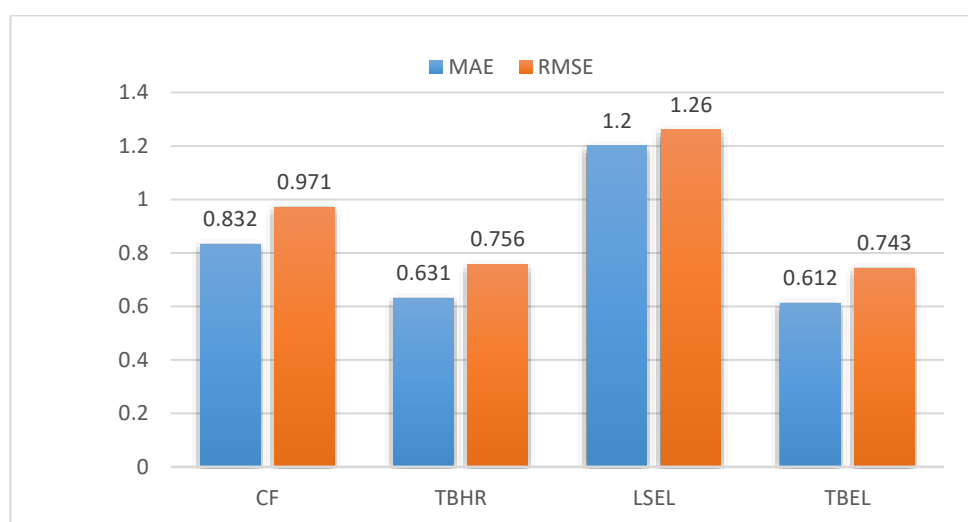
مقادیر عددی نتایج به دست آمده برای بررسی دقیق‌تر در جدول ۹ نیز ارائه شده است.

میزان بهبود پارامتر MAE در روش TBEL نسبت به روش‌های LSEL، TBHR و CF به ترتیب ۰/۹۶، ۰/۳۱ و ۰/۳۶ بوده است. همچنین در پارامتر RMSE نیز نسبت به روش‌های مذکور به ترتیب ۰/۶۹، ۰/۱۷ و ۰/۳۰ بهبود حاصل شده است. میزان بهبود RC در روش TBEL در مقایسه با روش‌های NPR_eL، TBHR و CF به ترتیب ۰/۱۰/۹، ۰/۰/۴۷ و ۰/۶۲ بوده است. همچنین در پارامتر UC نیز نسبت به روش‌های مذکور به ترتیب ۰/۱۴/۸، ۰/۳/۲ و ۰/۴۴/۲ بهبود حاصل شده است.

بنابراین نتایج مقایسه‌ها نشان می‌دهد با توجه به استفاده از اعتماد برای محاسبه شباهت نهایی کاربران، دقت بیشتری در ارائه پیشنهادها نسبت به روش‌های دیگر حاصل شده است. همچنین با توجه به نوع خوشه‌بندی فازی استفاده شده، مقدار پوشش رتبه بهتری برای کاربران و دوره‌های آموزشی نیز به دست آمده است.

است. TBEL به علت بررسی و محاسبه میزان اعتماد میان کاربران بر اساس افتخاراتی که در دوره‌های آموزشی پیشین به‌دست آورده‌اند و استفاده از آن در فرایند محاسبه شباهت میان کاربران خطای کمتری نسبت به دیگر روش‌ها داشته است. در واقع در فرایند تولید پیشنهاد برای کاربر هدف، با هدف ارائه پیشنهادها دقیق، از دوره‌های آموزشی پیشین کاربرانی استفاده می‌شود که در شبکه اعتماد کاربر هدف قرار دارند. مقادیر عددی نتایج به دست آمده برای بررسی دقیق‌تر در جدول ۸ نیز ارائه شده است.

شکل ۳ نشان‌دهنده میزان UC و RC روش پیشنهادی TBEL و سه روش دیگر است. نتایج نشان می‌دهد که روش پیشنهادی دارای پوشش بیشتری نسبت به سه روش TBHR، NPR_eL و CF است. TBEL به علت استفاده از روش خوشه‌بندی فازی و در نظر گرفتن هم‌پوشانی بیشتر خوشه‌ها، پوشش بیشتری نسبت به دیگر روش‌ها دارد. مقدار عددی نتایج به دست آمده برای بررسی دقیق‌تر در جدول ۹ نیز ارائه شده است.



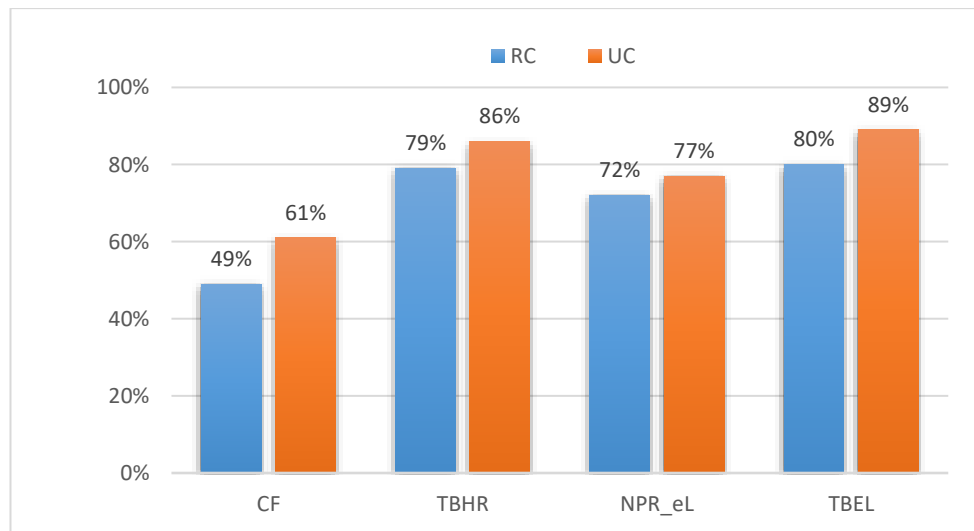
شکل ۲: مقایسه TBEL با روش‌های دیگر بر اساس معیارهای MAE و RMSE

Fig. 2: Comparison of TBEL with the other methods based on RMSE and MAE metrics

جدول ۸: نتایج TBEL و روش‌های دیگر بر اساس معیارهای MAE و RMSE

Table 8: Results of TBEL and the other methods based on RMSE and MAE metrics

معیارها Metrics	میانگین خطای مطلق MAE				خطای جذر میانگین مربعات RMSE			
	CF	TBHR	LSEL	TBEL	CF	TBHR	LSEL	TBEL
روش‌ها Methods								
نتایج Results	0.832	0.631	1.200	0.612	0.971	0.756	1.260	0.743



شکل ۳: مقایسه TBEL با روش‌های دیگر بر اساس معیارهای RC و UC

Fig. 3: Comparison of the TBEL and three other methods based on UC and RC measures in the Moodle dataset

جدول ۹: نتایج TBEL و روش‌های دیگر بر اساس معیارهای RC و UC

Table 9: Results of TBEL and the other methods based on UC and RC metrics

معیارها Metrics	نرخ پوشش رتبه‌ها RC (%)				نرخ پوشش کاربرها UC (%)			
	CF	TBHR	NPR_eL	TBEL	CF	TBHR	NPR_eL	TBEL
روش‌ها Methods								
نتایج Results	49.46	79.81	72.29	80.19	61.92	86.54	77.81	89.35

نتیجه‌گیری

را کاهش می‌دهد. بر اساس نتایج ارزیابی‌های انجام‌شده بر روی مجموعه داده Moodle، استفاده از روش‌های خوشه‌بندی فازی و قوانین انجمنی وزن‌دار در کنار یکدیگر در سیستم‌های پیشنهاددهنده یادگیری الکترونیکی مبتنی بر اعتماد، منجر به کاهش معیارهای خطای جذر میانگین مربعات و میانگین خطای مطلق و افزایش معیارهای پوشش رتبه و کاربر نسبت به سایر روش‌های مشابه شده است.

مشارکت نویسندگان

نویسندگان مقاله به طور فعال در ارائه ایده پژوهشی، طراحی معماری سیستم، جمع‌آوری داده و پیاده‌سازی و شبیه‌سازی آن مشارکت داشتند.

تشکر و قدردانی

نویسندگان این پژوهش از داوران محترم و دست‌اندرکاران نشریه فناوری آموزش کمال تشکر و قدردانی را دارند.

تعارض منافع

«هیچگونه تعارض منافع توسط نویسندگان بیان نشده است.»

در این مقاله، یک سیستم پیشنهاددهنده یادگیری الکترونیکی مبتنی بر اعتماد ارائه شده است که از روش خوشه‌بندی فازی و قوانین انجمنی وزن‌دار استفاده می‌کند. این روش مبتنی بر مدل است و بر اساس پالایش مشارکتی پیشنهادها را ارائه می‌دهد. برای غلبه بر خلوت بودن داده‌ها و بهبود فرایند یافتن کاربران مشابه با کاربر هدف، روش پیشنهادی بر اساس ماتریس‌های رتبه و اعتماد میان کاربران، یک شبکه اعتماد مبتنی بر روابط اعتماد بین کاربران برای کاربر هدف ایجاد می‌کند. در فرایند ارائه پیشنهاد دوره آموزشی به کاربر هدف از خوشه‌بندی فازی داده‌ها و اعمال قوانین انجمنی وزن‌دار بر روی اطلاعات کاربرانی استفاده می‌شود که در شبکه اعتماد کاربر هدف قرار دارند. به‌کارگیری هدف از به‌کارگیری معیار اعتماد میان کاربران سیستم‌های پیشنهاددهنده یادگیری الکترونیکی، افزایش دقت در انتخاب همسایه‌های کاربر هدف و محدود کردن اثرات مخرب کاربران و نظرات بی‌اعتبار آن‌ها است.

از طرف دیگر، با توجه به خوشه‌بندی فازی کاربران، پیش‌بینی رتبه دوره‌های آموزشی مختلف توسط روش پیشنهادی فقط بر اساس همسایه‌های موجود در خوشه کاربر هدف، انجام می‌شود که برای حجم انبوه اطلاعات، عملکرد بهتری خواهد داشت و مشکل خلوت بودن داده‌ها

منابع و مأخذ

system for e-learning forum. *Journal of Educational Technology & Society*. 2018;21 (1):112-125.

[15] Khanal SS, Prasad PW, Alsadoon A, Maag A. A systematic review: machine learning based recommendation systems for e-learning. *Education and Information Technologies*. 2019; ۲۵: 2635–2664

[16] Karga S, Satratzemi M. A hybrid recommender system integrated into LAMS for learning designers. *Education and Information Technologies*. 2018; 23 (3): 1297-1329.

[17] Aeiad E, Meziane F. An adaptable and personalised e-learning system applied to computer science programmes design. *Education and Information Technologies*. 2019; 24 (2):1485-1509.

[18] De Medio C, Limongelli C, Sciarrone F, Temperini M. MoodleREC: A recommendation system for creating courses using the moodle e-learning platform. *Computers in Human Behavior*. 2020; 104: 106168.

[19] Benhamdi S, Babouri A, Chiky R. Personalized recommender system for e-Learning environment. *Education and Information Technologies*. 2017; 22 (4):1455-1477.

[20] Parvin H, Moradi P, Esmaeili S. TCFACO: Trust-aware collaborative filtering method based on ant colony optimization. *Expert Systems with Applications*. 2019; 118: 152-168.

[21] Yadav S, Kumar V, Sinha S, Nagpal S. Trust aware recommender system using swarm intelligence. *Journal of computational science*. 2018; 28: 180-192.

[22] Xu K, Zhang W, Yan Z. A privacy-preserving mobile application recommender system based on trust evaluation. *Journal of computational science*. 2018; 26:87-107.

[23] Nobahari V, Jalali M, Mahdavi SJ. ISoTrustSeq: a social recommender system based on implicit interest, trust and sequential behaviors of users using matrix factorization. *Journal of Intelligent Information Systems*. 2019; 52 (2) :239-268.

[24] Jiang L, Cheng Y, Yang L, Li J, Yan H, Wang X. A trust-based collaborative filtering algorithm for e-commerce recommendation system. *Journal of Ambient Intelligence and Humanized Computing*. 2019; 10 (8):3023-3034.

[25] Pan Y, He F, Yu H, Li H. Learning adaptive trust strength with user roles of truster and trustee for trust-aware recommender systems. *Applied Intelligence*. 2020; 50 (2):314-327.

[26] Dwivedi P, Bharadwaj KK. Effective trust-aware e-learning recommender system based on learning styles and knowledge levels. *Journal of Educational Technology & Society*. 2013; 16 (4): 201-216.

[1] Bobadilla J, Ortega F, Hernando A, Gutiérrez A. Recommender systems survey. *Knowledge-based systems*. 2013; 46:109-132.

[2] Kunaver M, Požrl T. Diversity in recommender systems—A survey. *Knowledge-Based Systems*. 2017; 123 :154-162.

[3] Yang X, Guo Y, Liu Y, Steck H. A survey of collaborative filtering based social recommender systems. *Computer Communications*. 2014; 41 :1-10.

[4] Bhaskaran S, Santhi B. An efficient personalized trust- based hybrid recommendation (TBHR) strategy for e-learning system in cloud computing. *Cluster Computing*. 2019;22 (1):1137-1149.

[5] Nafea SM, Siewe F, He Y. A novel algorithm for course learning object recommendation based on student learning styles. In 19 Feb 2019 *International Conference on Innovative Trends in Computer Engineering (ITCE)*; IEEE; 2019 pp. 192-201.

[6] Chen S, Xu Z, Tang Y. A hybrid clustering algorithm based on fuzzy c-means and improved particle swarm optimization. *Arabian Journal for Science and Engineering*. 2014;39 (12): 8875-8887.

[7] Hassan T. Trust and trustworthiness in social recommender systems. In Companion Proceedings of the 2019 World Wide Web Conference; 2019 May 13. pp. 529-532.

[8] Klačnja-Milićević A, Vesin B, Ivanović M, Budimac Z. E-Learning personalization based on hybrid recommendation strategy and learning style identification. *Computers & Education*. 2011; 56 (3):885-899.

[9] Masud M. Collaborative e-learning systems using semantic data interoperability. *Computers in Human Behavior*. 2016; 61:127-135.

[10] Tarus JK, Niu Z, Kalui D. A hybrid recommender system for e-learning based on context awareness and sequential pattern mining. *Soft Computing*. 2018; 22 (8) :2449-2461.

[11] Dahdouh K, Dakkak A, Oughdir L, Ibriz A. Large-scale e-learning recommender system based on Spark and Hadoop. *Journal of Big Data*. 2019 ; 6 (2): 1-23.

[12] Klačnja-Milićević A, Ivanović M, Vesin B, Budimac Z. Enhancing e-learning systems with personalized recommendation based on collaborative tagging techniques. *Applied Intelligence*. 2018; 48 (6) :1519-1535.

[13] Wan S, Niu Z. An e-learning recommendation approach based on the self-organization of learning resource. *Knowledge-Based Systems*. 2018; 160: 71-87.

[14] Albatayneh NA, Ghauth KI, Chua FF. Utilizing learners' negative ratings in semantic content-based recommender

- [40] Xu K. (2013) Clustering. In: Dubitzky W., Wolkenhauer O., Cho KH., Yokota H. (eds.) *Encyclopedia of Systems Biology*. Springer, New York, NY
- [41] Soman KP, Diwakar S, Ajay V. Data mining: Theory and practice. PHI Learning Pvt. Ltd.; 2006.
- [42] Nayak J, Naik B, Behera HS. Fuzzy C-means (FCM) clustering algorithm: A decade review from 2000 to 2014. In *Computational intelligence in data mining-volume 2*. New Delhi: Springer; 2015. p. 133-149.
- [43] Stetco A, Zeng XJ, Keane J. Fuzzy C-means++: Fuzzy C-means with effective seeding initialization. *Expert Systems with Applications*. 2015; 42 (21):7541-7548.
- [44] Pelleg D, Moore AW. X-means: Extending k-means with efficient estimation of the number of clusters. In *ICML 2000 Jun 29*. pp. 727-734
- [45] Saleem F, Iltaf N, Afzal H, Shahzad M. Using trust in collaborative filtering for recommendations. In *IEEE 28th International Conference on Enabling Technologies: Infrastructure for Collaborative Enterprises (WETICE)*: IEEE; 2019 Jun 12. pp. 214-222.
- [46] Moradi P, Ahmadian S. A reliability-based recommendation method to improve trust-aware recommender systems. *Expert Systems with Applications*. 2015; 42 (21):7386-7398.
- [47] Ardissono L, Mauro N. A compositional model of multi-faceted trust for personalized item recommendation. *Expert Systems with Applications*. 2020; 140:112880.
- [48] Shchetinin EY. On improving the reliability of recommender systems with users clustering. In *ITMM (Selected Papers) 2019*. pp. 118-129.
- [49] Yuan W, Guan D, Lee YK, Lee S. The small-world trust network. *Applied Intelligence*. 2011; 35 (3) :399-410.
- [50] Agrawal R, Srikant R. Fast algorithms for mining association rules. In *Proc. of the 20th VLDB Conference 1994 Sep 12*. pp. 487-499.
- [27] Deng X, Li H, Huangfu F. A trust-aware neural collaborative filtering for e-learning recommendation. *Educational Sciences: Theory & Practice*. 2018; 18 (5).
- [28] Wan S, Niu Z. A hybrid e-learning recommendation approach based on learners' influence propagation. *IEEE Transactions on Knowledge and Data Engineering*. 2019; 32 (5): 827-840.
- [29] Liu X. A collaborative filtering recommendation algorithm based on the influence sets of e-learning group's behavior. *Cluster Computing*. 2019; 22 (2):2823-2833.
- [30] Hinz VT, Pimenta MS. Integrating Reputation to Recommendation Techniques in an e-learning Environment. In *2018 17th International Conference on Information Technology Based Higher Education and Training (ITHET)* IEEE; 2018 Apr 26 pp. 1-6.
- [31] Hasan M, Roy F. An item-item collaborative filtering recommender system using trust and genre to address the cold-start problem. *Big Data and Cognitive Computing*. 2019;3 (3) :39.
- [32] Jiang L, Cheng Y, Yang L, Li J, Yan H, Wang X. A trust-based collaborative filtering algorithm for e-commerce recommendation system. *Journal of Ambient Intelligence and Humanized Computing*. 2019; 10 (8):3023-3034.
- [33] Si M, Li Q. Shilling attacks against collaborative recommender systems: a review. *Artificial Intelligence Review*. 2020 ; 53 (1):291-319.
- [34] Alonso S, Bobadilla J, Ortega F, Moya R. Robust model-based reliability approach to tackle shilling attacks in collaborative filtering recommender systems. *IEEE Access*. 2019; 7: 41782-41798.
- [35] Aggarwal CC. Attack-resistant recommender systems. In *Recommender Systems 2016*; 385-410. Springer, Cham.
- [36] Massa P, Avesani P. Trust metrics in recommender systems. In *Computing with Social Trust*; 2009 (pp. 259-285). Springer, London.
- [37] Victor P, Cornelis C, De Cock M, Teredesai A. Trust-and distrust-based recommendations for controversial reviews. In *Web Science Conference (WebSci'09: Society On-Line)* 2009 (No. 161).
- [38] Zuo X, Wei X, Yang B. Trust-distrust aware recommendation by integrating metric learning with matrix factorization. In *International Conference on Knowledge Science, Engineering and Management*: 2018 Aug 17 (pp. 361-370). Springer, Cham.
- [39] Oh HK, Kim JW, Kim SW, Lee K. A unified framework of trust prediction based on message passing. *Cluster Computing*. 2019; 22 (1):2049-2061.

معرفی نویسندگان

AUTHOR(S) BIOSKETCHES



رضوان محمدرضایی دانشجوی دکتری سیستم‌های نرم‌افزاری در واحد تهران مرکزی - دانشگاه آزاد اسلامی می‌باشند. مدرک کارشناسی مهندسی کامپیوتر نرم‌افزار را در سال ۱۳۸۶ از دانشگاه آزاد اسلامی واحد ماهشهر و مدرک کارشناسی ارشد



تخصصی مهندسی کامپیوتر را به ترتیب در سال‌های ۱۳۷۸ و ۱۳۸۴ از واحد علوم و تحقیقات - دانشگاه آزاد اسلامی دریافت کرده‌اند. ایشان تا کنون بیش از ۵۰ مقاله علمی در مجلات و همایش‌های معتبر بین‌المللی به چاپ رسانده‌اند.

زمینه‌های تخصصی ایشان عبارتند از: تحلیل و مدیریت داده‌های کلان، سیستم‌های توزیع شده، تحلیل‌های مرتبط با شبکه‌های اجتماعی.

Ravanmehr, R. Assistant Professor, Department of Computer Engineering, Central Tehran Branch, Islamic Azad University, Tehran, Iran

r.ravanmehr@iauctb.ac.ir

مهندسی کامپیوتر - نرم‌افزار را در سال ۱۳۹۰ از دانشگاه علوم و تحقیقات اهواز دریافت نمود. تا کنون چندین مقاله علمی در مجلات و همایش‌های بین‌المللی توسط ایشان ارائه و چاپ شده است. سیستم‌های پیشنهاددهنده، یادگیری ماشین و تحلیل داده‌ها زمینه‌های تحقیقاتی مورد علاقه ایشان هستند.

Mohamadrezai, R. PhD Student, Department of Computer Engineering, Central Tehran Branch, Islamic Azad University, Tehran, Iran

rez.mohamadrezaeilarki.eng@iauctb.ac.ir

رضا روانمهر استادیار مهندسی کامپیوتر در واحد تهران مرکزی - دانشگاه آزاد اسلامی می‌باشند. ایشان مدرک کارشناسی مهندسی کامپیوتر را از دانشگاه شهید بهشتی و مدارک کارشناسی ارشد و دکتری

Citation (Vancouver): Mohamadrezai R, Ravanmehr R. [A trust-based recommender system for e-Learning environment using fuzzy clustering]. *Tech. Edu. J.* 2021; 15(3): 439-464

 <http://dx.doi.org/10.22061/tej.2021.6807.2454>



COPYRIGHTS

©2021 The author(s). This is an open access article distributed under the terms of the Creative Commons Attribution (CC BY 4.0), which permits unrestricted use, distribution, and reproduction in any medium, as long as the original authors and source are cited. No permission is required from the authors or the publishers.