

Summer 2024, 5 (2), 119-135
DOR:

Received: 28 Apr 2024
Accepted: 9 June 2024

مقاله پژوهشی

Enhancing Adaptive Learning Systems: A Probabilistic Approach Using Bayesian Networks and User Clustering

Parvaneh Ahmadifard¹, Mojtaba Salehi^{2*}, Mohamad Yarahmadi³, Morteza Zakeri Nasr Abadi⁴

1. MSc, Department of Computer Engineering, Khorramabad Branch, Islamic Azad University, Khorramabad, Iran.
2. Instructor, Department of Computer Engineering, Khorramabad Branch, Islamic Azad University, Khorramabad, Iran.
**Corresponding Author, Mo.Salehi@iau.ac.ir*
3. Instructor, Department of Math, Khorramabad Branch, Islamic Azad University, Khorramabad, Iran.
4. Assistant Professor, School of Computer Engineering, Amirkabir University of Technology (Tehran Polytechnic), Tehran, Iran.

Abstract

Introduction: This study proposes an innovative framework to enhance the quality and adaptability of web-based educational systems using Bayesian networks. The proposed methodology aims to address issues like the cold-start problem and improve the scalability and personalization of educational services.

Method: The proposed methodology encompasses feature selection through the Relief algorithm, user clustering via k-means, and probabilistic service recommendation based on Bayesian inference. Feature selection prioritizes user-specific attributes, reducing data dimensionality and enhancing computational efficiency. Clustering organizes users into homogeneous groups based on similarities, facilitating tailored service provision. The system was validated using a dataset of 2000 users, each characterized by 16 features and linked to 24 services.

Results: The experimental results demonstrated an accuracy rate of 88.15%, underscoring the model's effectiveness compared to traditional techniques. Bayesian networks enabled the analysis of probabilistic relationships and uncertainties, ensuring that the most relevant services were recommended to each user.

Discussion: The proposed approach addressed the cold-start problem, improving scalability and adaptability to diverse user needs. Key advantages of this framework include enhanced personalization, computational efficiency, and significant improvements in user satisfaction. The study's findings highlight the potential of intelligent learning systems to adjust dynamically to user requirements, paving the way for more engaging and effective educational experiences. Future research will incorporate heuristic algorithms to improve prediction accuracy and integrate real-time user feedback to refine the recommendation process further.

Keywords: Intelligent tutoring system, feature selection, Bayesian network, K-means algorithm.

بهبود سیستم‌های یادگیری تطبیقی: یک رویکرد احتمالی با استفاده از شبکه‌های بیزین و خوشه‌بندی کاربر

دوره پنجم، تابستان ۱۴۰۳
شماره دوم، صص: ۱۱۹-۱۳۵

تاریخ دریافت: ۱۴۰۳/۰۲/۰۹
تاریخ پذیرش: ۱۴۰۳/۰۳/۲۰

پروانه احمدی فرد^۱، مجتبی صالحی^{۲*}، محمد یاراحمدی^۳، مرتضی ذاکری^۴

- ۱- کارشناسی ارشد، گروه کامپیوتر، واحد خرم آباد، دانشگاه آزاد اسلامی، خرم آباد، ایران. ahmadicom2004@yahoo.com
- ۲- مربی، گروه کامپیوتر، واحد خرم آباد، دانشگاه آزاد اسلامی، خرم آباد، ایران. (نویسنده مسئول) mo.salehi@iau.ac.ir
- ۳- مربی، گروه ریاضی، واحد خرم آباد، دانشگاه آزاد اسلامی، خرم آباد، ایران. mohammadyarahmadi.1981@iau.ac.ir
- ۴- استادیار، دانشکده مهندسی کامپیوتر، دانشگاه صنعتی امیرکبیر (پلی تکنیک تهران)، تهران، ایران. zakeri@aut.ac.ir

چکیده: این پژوهش یک چارچوب نوآورانه برای ارتقای کیفیت و انطباق‌پذیری سامانه‌های آموزشی مبتنی بر وب با استفاده از شبکه‌های بیزین ارائه می‌دهد. روش پیشنهادی شامل انتخاب ویژگی از طریق الگوریتم Relief، خوشه‌بندی کاربران با استفاده از الگوریتم k-means و پیشنهاد خدمات به صورت احتمالاتی بر اساس استنتاج بیزین است. انتخاب ویژگی با اولویت‌بندی ویژگی‌های کلیدی کاربران، ابعاد داده‌ها را کاهش داده و کارایی پردازشی را افزایش می‌دهد. خوشه‌بندی، کاربران را بر اساس شباهت‌هایشان در گروه‌های همگن سازمان‌دهی کرده و ارائه خدمات شخصی‌سازی شده را تسهیل می‌کند. این سامانه با استفاده از یک مجموعه داده شامل ۲۰۰۰ کاربر که هر کدام دارای ۱۶ ویژگی و ۲۴ خدمت بودند، ارزیابی شد. شبکه‌های بیزین روابط احتمالاتی و عدم قطعیت‌ها را تحلیل کرده و خدمات مرتبط را به هر کاربر پیشنهاد می‌دهند. این روش مشکل آغاز سرد را برطرف کرده و انطباق‌پذیری و مقیاس‌پذیری را برای نیازهای متنوع کاربران بهبود می‌بخشد. نتایج آزمایش‌ها دقتی معادل ۸۸٫۱۵ درصد را نشان داد که مؤثر بودن مدل در مقایسه با روش‌های سنتی را تأیید می‌کند. مزایای کلیدی شامل بهبود شخصی‌سازی، افزایش کارایی محاسباتی و ارتقای رضایت کاربران است. یافته‌های این مطالعه، توانایی سامانه‌های آموزشی هوشمند را برای تطبیق پویا با نیازهای کاربران نشان داده و زمینه‌ساز تجربه‌های آموزشی جذاب‌تر و مؤثرتر می‌شود. پژوهش‌های آتی بر ادغام الگوریتم‌های فراابتکاری برای پیش‌بینی دقیق‌تر و دریافت بازخورد بلندمدت کاربران به منظور بهبود فرآیند پیشنهاد متمرکز خواهند بود.

واژه‌های کلیدی: سیستم آموزش هوشمند، انتخاب ویژگی، شبکه بیزین، الگوریتم K-Means.

۱. مقدمه

تحولات فناوری اطلاعات و ارتباطات فرصت‌های بی‌سابقه‌ای را برای گسترش دسترسی به آموزش ایجاد کرده‌است. آموزش الکترونیکی با ارائه خدمات آموزشی انعطاف‌پذیر و بدون محدودیت زمانی و مکانی، نقش مهمی در تحقق هدف آموزش برای همه ایفا می‌کند. با این حال، اثربخشی این سامانه‌ها به شدت به توانایی آن‌ها در ارائه خدمات شخصی‌سازی شده و متناسب با نیازهای کاربران وابسته است. شخصی‌سازی می‌تواند تجربه یادگیری کاربران را بهبود بخشد و در نهایت به افزایش کارایی سامانه‌ها و نتایج آموزشی منجر شود [۱].

در سال‌های اخیر، روش‌های هوش مصنوعی، به‌ویژه شبکه‌های بیزین، به‌عنوان ابزاری کارآمد برای تحلیل داده‌های آموزشی و مدل‌سازی روابط پیچیده مورد توجه قرار گرفته‌اند. شبکه‌های بیزین به دلیل قابلیت استنتاج احتمالاتی و مدیریت داده‌های ناقص، ابزار مناسبی برای غلبه بر محدودیت‌های روش‌های سنتی داده‌کاوی محسوب می‌شوند [۲]. پژوهش‌هایی همچون استفاده از شبکه‌های بیزین برای بهینه‌سازی سیستم‌های آموزشی مبتنی بر وب و ایجاد پلتفرم‌های یادگیری شخصی‌سازی شده، نشان داده‌اند که این مدل‌ها می‌توانند با ارزیابی بلادرنگ توانایی‌های یادگیرندگان و تحلیل روابط دانش، کارایی یادگیری را به‌طور چشمگیری بهبود بخشند [۳] [۴].

همچنین، رویکردهای مبتنی بر شبکه‌های بیزین به‌عنوان ابزارهای کلیدی برای توسعه سیستم‌های آموزشی تطبیقی مطرح شده‌اند که امکان ارائه محتوای آموزشی متناسب با مهارت و نیازهای کاربران را فراهم می‌کنند [۵]. با وجود این، اغلب پژوهش‌ها بر جنبه‌های نظری متمرکز بوده و نیاز به مطالعات عملیاتی و ترکیبی برای بهره‌گیری کامل از این رویکردها احساس می‌شود.

سامانه‌های آموزش الکترونیکی اگرچه مزایای بسیاری دارند، اما همچنان با چالش‌های متعددی مواجه‌اند. یکی از چالش‌های مهم این سیستم‌ها، عدم توانایی کافی در ارائه خدمات شخصی‌سازی شده و متناسب با نیازهای کاربران است. این مسئله نه تنها بر تجربه یادگیری کاربران تأثیر منفی می‌گذارد، بلکه کارایی کلی سامانه‌ها را نیز کاهش می‌دهد [۶].

یکی دیگر از مشکلات عمده، مسئله آغاز سرد است که در آن سیستم نمی‌تواند برای کاربران جدید، خدمات مناسبی پیشنهاد کند زیرا داده‌های کافی برای تحلیل رفتار آن‌ها وجود ندارد. این مشکل باعث می‌شود که تجربه کاربران در مراحل اولیه استفاده از سیستم کاهش یابد و انگیزه ادامه یادگیری تحت تأثیر قرار گیرد [۷].

علاوه بر این، سیستم‌های آموزش الکترونیکی معمولاً در مواجهه با داده‌های ناقص یا ناسازگار عملکرد مطلوبی ندارند. نبود چارچوب‌هایی که بتوانند این مشکلات را مدیریت کرده و از ابزارهای پیشرفته‌ای مانند شبکه‌های بیزین بهره‌گیرند، منجر به محدودیت در استفاده عملی از این سیستم‌ها شده‌است. به‌طور خاص، فقدان ترکیب شبکه‌های بیزین با تکنیک‌های داده‌کاوی نظیر خوشه‌بندی و انتخاب ویژگی،

مانع از بهره‌برداری کامل از ظرفیت‌های این فناوری‌ها شده‌است [۸]. همچنین، بررسی تأثیر این رویکردها در محیط‌های واقعی که داده‌های متنوع و غیرقابل پیش‌بینی دارند، به اندازه کافی انجام نشده‌است. بیشتر پژوهش‌ها در محیط‌های آزمایشگاهی با داده‌های محدود صورت گرفته و نمی‌توانند به‌طور کامل پیچیدگی‌های واقعی را منعکس کنند [۹].

پژوهش‌های پیشین عمدتاً به اثربخشی شبکه‌های بیزین در پیش‌بینی رفتار کاربران و تحلیل داده‌های آموزشی پرداخته‌اند، اما به‌ندرت به ارائه چارچوب‌های عملیاتی و ترکیب این مدل‌ها با روش‌های دیگر مانند خوشه‌بندی و انتخاب ویژگی پرداخته شده‌است. به‌عنوان مثال، در یک مطالعه، استفاده از شبکه‌های بیزین در سیستم‌های OER تطبیقی به بهبود تسلط بر مهارت‌ها و افزایش نتایج یادگیری دانش‌آموزان کمک کرده‌است [۱۰]. همچنین، استفاده از شبکه‌های بیزین برای ارزیابی نتایج یادگیری الکترونیکی نشان داده‌است که این مدل‌ها می‌توانند اثربخشی آموزش آنلاین را با تحلیل داده‌های یادگیرندگان افزایش دهند [۱۱].

در زمینه بومی نیز، محدودیت‌های زیرساختی، مشکلات دسترسی به داده‌های کافی، و نیاز به انطباق با شرایط فرهنگی و آموزشی محلی، چالش‌های جدی برای توسعه این سامانه‌ها ایجاد کرده‌است. این شکاف‌ها نشان می‌دهند که نیاز به ارائه راهکارهای ترکیبی و عملیاتی برای بهره‌گیری بهتر از شبکه‌های بیزین در سامانه‌های آموزش الکترونیکی وجود دارد [۱۲].

مطالعات متعدد اثربخشی شبکه‌های بیزین در بهبود عملکرد سیستم‌های آموزشی را نشان داده‌اند. به‌عنوان مثال، در یک مطالعه استفاده از شبکه‌های بیزین برای تحلیل داده‌های یادگیرندگان به کاهش خطاهای پیش‌بینی و بهبود دقت تصمیم‌گیری منجر شد [۱۳]. علاوه بر این، ترکیب شبکه‌های بیزین با الگوریتم‌های خوشه‌بندی، امکان ارائه خدمات تطبیقی متناسب با گروه‌های کاربری مختلف را فراهم کرده‌است [۱۴].

در ایران، سامانه‌های آموزش الکترونیکی با محدودیت‌هایی مانند عدم دسترسی همگانی به اینترنت، نبود منابع کافی و مهارت اندک کاربران در استفاده از فناوری مواجهند. این چالش‌ها، نیاز به ارائه چارچوب‌های بومی‌سازی شده را دوچندان کرده‌است.

این مطالعه با هدف طراحی یک چارچوب جامع برای بهینه‌سازی سامانه‌های آموزش الکترونیکی مبتنی بر وب انجام شده‌است. اهداف خاص پژوهش عبارتند از:

- ارائه یک چارچوب ترکیبی شامل شبکه‌های بیزین، خوشه‌بندی و انتخاب ویژگی.
- بهبود دقت پیش‌بینی و شخصی‌سازی خدمات آموزشی.
- بررسی عملی اثربخشی چارچوب پیشنهادی در سناریوهای واقعی.

این پژوهش تلاش می‌کند تا با رفع شکاف‌های موجود، زمینه‌ساز توسعه سامانه‌های آموزش الکترونیکی کارآمد و انطباق‌پذیر در مقیاس

محلی و جهانی شود. در این مقاله با استفاده از شبکه‌های بیزین یک چارچوب نوآورانه برای ارتقای کیفیت سامانه‌های آموزشی مبتنی بر وب آورده شده، با استفاده انتخاب ویژگی از طریق الگوریتم Relief و خوشه‌بندی کاربران با استفاده از الگوریتم k-means، ابعاد داده‌ها را کاهش داده و کارایی پردازش افزایش یافته، این مزیت یکی از دلایل ارائه این روش پیشنهاد می‌باشد. در بخش دوم مروری بر کارهای پیشین مرتبط با این تحقیق بیان شده‌است. در بخش سوم روش پیشنهادی بیان شده‌است. در این بخش ابتدا به‌طور کلی مراحل انجام کار شرح داده شده‌است سپس انتخاب ویژگی و خوشه‌بندی بیان شده‌است. در بخش چهارم پیشنهاد سرویس به کاربران بیان شده‌است. در بخش پنجم آزمایش‌های همراه با تحلیل نتایج ذکر شده‌است. در بخش ششم معیارهای ارزیابی مورد بررسی قرار گرفته‌اند و سپس نهایتاً در بخش آخر پژوهش نتیجه‌گیری و پیشنهاد کارهای آتی برای ادامه این پژوهش بیان شده‌است.

۲. کارهای مرتبط

مدل فراگیر در مقاله [۱۵] به‌عنوان نمایشی از باورهای یک سیستم رایانه‌ای در مورد فراگیران تعریف شده‌است. به عبارت دیگر، این مدل یک نمایش انتزاعی از ویژگی‌ها و رفتارهای احتمالی یک فراگیر را ارائه می‌دهد. مدل‌سازی فراگیر شامل مجموعه‌ای از طرز فکرها و برخوردهایی است که یک فرد می‌تواند در طول فرایند یادگیری از خود نشان دهد. این مدل‌ها درواقع نقش مهمی در شبیه‌سازی فرایندهای یادگیری و تحلیل تعاملات مختلف فراگیران ایفا می‌کنند. با این حال، باید بین مدل فراگیر و مدل کاربر که جنبه‌ای جامع‌تر دارد و به رفتارهای عمومی‌تری می‌پردازد که فقط به تدریس محدود نمی‌شود، تمایز قائل شد. مدل کاربر معمولاً در چارچوب وسیع‌تری از تعاملات انسان با سیستم به‌کار می‌رود، در حالی که مدل فراگیر به‌طور خاص بر جنبه‌های یادگیری و نیازهای آموزشی متمرکز است. در سیستم‌های آموزشی هوشمند (ITS) فراگیران اجزای اصلی سیستم هستند و بر اساس مدل‌های پیچیده‌تری از رفتار و یادگیری آن‌ها طراحی می‌شوند [۱۵].

مدل‌های فراگیر در چارچوب نظریه‌های مختلف یادگیری مانند رفتارگرایی، شناخت‌گرایی و ساخت‌گرایی مورد بررسی قرار می‌گیرند. در رویکرد رفتارگرایی، یادگیری به‌عنوان یک فرآیند واکنشی به محرک‌ها دیده می‌شود، جایی که فرد مانند یک جعبه سیاه عمل می‌کند و پاسخ‌هایش به تقویت محرک‌ها بستگی دارد. این رویکرد بیشتر به تکرار و تقویت پاسخ‌های درست تأکید دارد. در مقابل، نظریه شناخت‌گرایی بر پردازش اطلاعات و اصلاح نمایش‌های ذهنی تأکید می‌کند و یادگیری را به‌عنوان تغییرات در این نمایش‌ها می‌بیند. همچنین، نظریه ساخت‌گرایی فرض می‌کند که فراگیران دانش را بر اساس تجربیات گذشته خود می‌سازند و هر فرد واقعیت و دانش خود را ایجاد می‌کند. این رویکرد تأکید زیادی بر یادگیری فعال و تجربیات شخصی دارد. در این میان،

مدل‌های ITS سنتی بیشتر از رفتارگرایی و شناخت‌گرایی حمایت می‌کنند، چراکه بیشتر به ارزیابی مهارت‌ها و شبیه‌سازی وضعیت‌های درونی فراگیران توجه دارند. با این وجود، به‌کارگیری مدل‌های پیچیده‌تر و پیشرفته‌تر، مانند استفاده از شبکه‌های بیزین، می‌تواند به شکل‌دهی مدل‌هایی کمک کند که نه تنها به ارزیابی مهارت‌ها، بلکه به شبیه‌سازی فرایندهای پیچیده‌تر یادگیری و تعاملات فراگیران می‌پردازند. انواع مختلفی از مدل‌های فراگیر وجود دارد که در مطالعات متعدد به آن‌ها پرداخته شده‌است. این مدل‌ها شامل مدل‌های پوششی، مدل‌های افتراقی و مدل‌های اختلالی هستند [۱۷][۱۸][۱۹][۲۰]. هر یک از این مدل‌ها به‌طور خاص بر جنبه‌های متفاوت یادگیری و رفتارهای فردی متمرکز هستند و می‌توانند به بهبود دقت سیستم‌های آموزشی هوشمند کمک کنند. در این راستا، مدل‌های شبکه بیزین نیز به‌عنوان ابزارهایی قدرتمند برای تحلیل و پیش‌بینی رفتارهای یادگیری دانش‌آموزان و طراحی سیستم‌های آموزشی شخصی‌سازی شده معرفی شده‌اند.

در مطالعه [۲۱] نویسندگان به بررسی استفاده از درخت تصمیم و شبکه عصبی برای شناسایی عوامل مؤثر در هدایت تحصیلی دانش‌آموزان پرداخته‌اند. این سیستم با تحلیل داده‌های وضعیت تحصیلی و مشاوره دانش‌آموزان، شانس موفقیت آن‌ها را پیش‌بینی کرده و به هدایت تحصیلی بهینه کمک می‌کند. همچنین، در پژوهش [۲۲]، به ارتقای کیفیت آموزش در سیستم‌های آموزش الکترونیکی با استفاده از داده‌کاوی آموزشی پرداخته شده‌است. این تحقیق هدف اصلی خود را کشف الگوهای نهفته در انتخاب واحدهای درسی و پیش‌بینی نمرات دانش‌آموزان قرار داده‌است. همچنین، تأثیر عواملی مانند ساعت و زمان ورود دانش‌آموزان به سیستم آموزش الکترونیکی بر کیفیت و نتایج یادگیری مورد بررسی قرار گرفته‌است. نتایج این تحقیق نشان می‌دهد که استفاده از تکنیک‌های داده‌کاوی می‌تواند به بهبود کیفیت آموزش و هدایت تحصیلی دانش‌آموزان کمک کند.

در پژوهش دیگری که در [۲۳] انجام شده‌است، نویسندگان به شناسایی عوامل مؤثر بر افت تحصیلی دانش‌آموزان با استفاده از قوانین انجمنی و تحلیل خوشه‌ای پرداخته‌اند. این تحقیق به دنبال پیاده‌سازی مدل‌های داده‌کاوی پیش‌بینی‌کننده برای پیش‌بینی وضعیت تحصیلی دانش‌آموزان بر اساس مشخصات فردی و تحصیلی آن‌ها بوده‌است. نتایج این تحقیق نشان داده‌است که با استفاده از این مدل‌ها می‌توان عوامل مؤثر بر افت تحصیلی را شناسایی کرده و از آن‌ها برای پیشگیری از افت تحصیلی و بهبود کیفیت ارتباط مسئولین و والدین با دانش‌آموزان بهره‌برد.

در پژوهش [۲۴] با استفاده از الگوریتم‌های مختلفی مانند درخت تصمیم بهبود یافته و الگوریتم نزدیک‌ترین همسایه، سوابق تحصیلی دانش‌آموزان را تجزیه و تحلیل کرده و به هدایت آن‌ها به رشته‌های تحصیلی مناسب کمک می‌کند. این تحقیق نشان داده‌است که ترکیب تکنیک‌های مختلف داده‌کاوی می‌تواند به‌طور مؤثری در هدایت تحصیلی دانش‌آموزان در مراحل مختلف تحصیلی مورد استفاده قرار گیرد.

۳. روش پیشنهادی

۱.۳ زمینه

روش پیشنهادی این پژوهش به منظور بهبود کیفیت خدمات در سیستم‌های آموزش الکترونیکی، بر پایه استفاده از شبکه‌های بیزین طراحی شده است. همان‌طور که در شکل (۱) نمایش داده شده، این روش شامل چند مرحله اصلی است که به‌طور خلاصه به شرح زیر توضیح داده می‌شود.

در اولین مرحله، مجموعه‌ای از ویژگی‌های کلیدی شناسایی می‌شوند که بیشترین تأثیر را در بهبود کیفیت خدمات به کاربران دارند. این ویژگی‌ها شامل عواملی هستند که نیازها، ترجیحات و الگوهای رفتاری کاربران در استفاده از خدمات آموزش الکترونیکی را نمایان می‌کنند. ویژگی‌هایی که تأثیر کمی بر کیفیت خدمات دارند یا تکراری و غیرضروری هستند، در این مرحله حذف می‌شوند. فرآیند انتخاب ویژگی‌ها نه تنها به ساده‌سازی مدل کمک می‌کند، بلکه حجم داده‌های پردازش شده را کاهش می‌دهد و در نتیجه، کارایی سیستم بهبود می‌یابد. این امر موجب می‌شود مدل سریع‌تر و مؤثرتر عمل کند.

در مرحله بعد، کاربران با استفاده از الگوریتم خوشه‌بندی k -means بر مبنای ویژگی‌های انتخاب شده گروه‌بندی می‌شوند. الگوریتم k -means یک روش شناخته شده است که داده‌ها را به خوشه‌های مختلف تقسیم می‌کند، به گونه‌ای که کاربران در هر خوشه بیشترین شباهت را به یکدیگر دارند و تفاوت‌ها بین خوشه‌ها به حداکثر می‌رسد. این فرآیند به سیستم این امکان را می‌دهد که کاربران را بر اساس رفتار و نیازهای مشابه دسته‌بندی کند. به این ترتیب، نیازها و الگوهای رفتاری کاربران در هر خوشه به‌طور دقیق‌تر شناسایی می‌شود و امکان ارائه خدمات اختصاصی‌تر به هر گروه فراهم می‌آید.

پس از خوشه‌بندی، سیستم به‌طور هوشمندانه سرویس‌ها و محتوای آموزشی را که بیشترین کاربرد را برای هر گروه از کاربران دارد، شناسایی و به آن‌ها پیشنهاد می‌دهد. این فرآیند موجب بهینه‌سازی تجربه کاربری می‌شود، زیرا کاربران هر گروه دقیقاً خدمات و محتوایی را دریافت می‌کنند که برای آن‌ها بیشترین فایده را دارد. به عنوان مثال، به کاربرانی که در زمینه خاصی نیاز به تقویت دارند، سرویس‌هایی پیشنهاد می‌شود که نقاط ضعف آن‌ها را پوشش دهد، در حالی که کاربران پیشرفته‌تر به محتوای سطح بالاتر و چالش‌برانگیزتر دسترسی پیدا می‌کنند.

در مرحله بعدی، شبکه‌های بیزین به عنوان ابزار اصلی برای تحلیل احتمالی و استنتاج به کار گرفته می‌شوند. شبکه بیزین یک مدل گرافیکی جهت‌دار است که روابط و وابستگی‌ها میان متغیرها را به صورت تصویری نمایش می‌دهد. این شبکه به سیستم این امکان را می‌دهد که حتی با داده‌های ناقص یا نامطمئن، پیش‌بینی‌هایی در خصوص رفتار و نیازهای آینده کاربران انجام دهد.

در نهایت، عملکرد سیستم در ارائه خدمات به کاربران مورد ارزیابی قرار می‌گیرد. نتایج حاصل از پیشنهاد سرویس‌ها و سطح رضایتمندی

در سال‌های اخیر، یادگیری الکترونیکی به عنوان یک روش آموزشی نوین مورد توجه قرار گرفته است. یکی از چالش‌های اصلی این حوزه، عدم تعامل مستقیم بین یادگیرنده و آموزش‌دهنده است که نیاز به شخصی‌سازی آموزش را افزایش می‌دهد. در این راستا، قره‌باغی و امیری در [۲۵] از الگوریتم شبه‌نظارتی LP-MLTSVM برای شناسایی سطح دانش یادگیرندگان و بهبود کیفیت آموزش استفاده کرده‌اند. نتایج نشان داد که این روش در مقایسه با روش‌های سنتی مانند درخت تصمیم و ماشین بردار پشتیبان، عملکرد و رضایت یادگیرندگان را به‌طور معناداری افزایش می‌دهد.

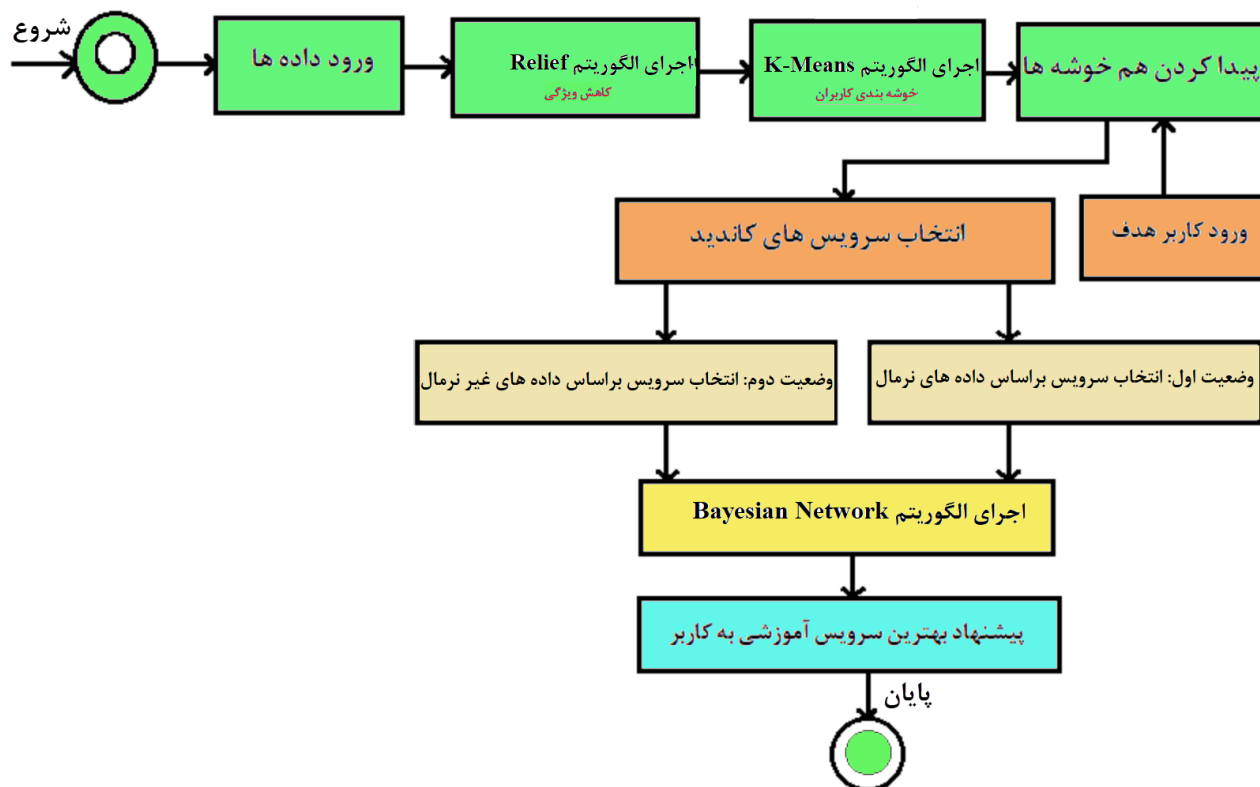
فارغ‌زاده و مدیری یک زیرساخت ابری آگاه از کارایی برای سیستم‌های آموزش الکترونیکی ارائه کردند که به صورت یکپارچه و پویا عملکرد کل پشته ابری را پایش می‌کند. این زیرساخت با استفاده از واحدهایی مانند مخزن دانش اشتراکی و واحد کنترل و اجرا، معیارهای کارایی مانند زمان پاسخ و گذردهی را بهبود بخشید. این رویکرد مستقل از مکانیزم‌های اجرایی خاص، قابلیت تطبیق با محیط‌های چندمستأجری را فراهم می‌کند [۲۶].

در پژوهشی دیگر [۲۷] با هدف بهبود دقت و شخصی‌سازی پیشنهادات فیلم، از ترکیب تکنیک‌های یادگیری عمیق و تحلیل معنایی استفاده می‌کند. روش پیشنهادی این مقاله، با بهره‌گیری از داده‌های کاربران و ویژگی‌های محتوایی فیلم‌ها، مدلی ارائه می‌دهد که قادر است ترجیحات گروهی و فردی کاربران را به طور همزمان در نظر بگیرد. نتایج آزمایش‌ها نشان می‌دهد که این رویکرد در مقایسه با روش‌های سنتی، عملکرد بهتری در پیش‌بینی ترجیحات کاربران و ارائه پیشنهادات مرتبط‌تر دارد. این مطالعه به عنوان یک چارچوب نوآورانه در حوزه سیستم‌های توصیه‌گر، به ویژه در زمینه توصیه‌های گروهی و معنایی، مطرح شده است.

در مجموع، استفاده از شبکه‌های بیزین و تکنیک‌های داده‌کاوی در تحلیل و مدل‌سازی سیستم‌های آموزشی، باعث پیش‌بینی دقیق‌تر وضعیت تحصیلی دانش‌آموزان و بهبود فرایندهای یادگیری آنلاین می‌شود. این روش‌ها می‌توانند اطلاعات پنهان و پیچیده در داده‌های آموزشی را شناسایی کرده و سیستم‌های آموزشی را به گونه‌ای بهینه‌سازی کنند که نتایج یادگیری فردی هر دانش‌آموز را بهبود دهند. همچنین، با توجه به پیچیدگی‌های موجود در فرآیند یادگیری و نیاز به تحلیل داده‌های عظیم، استفاده از مدل‌های پیشرفته مانند شبکه‌های بیزین می‌تواند به عنوان ابزاری مؤثر در طراحی سیستم‌های آموزشی هوشمند و شخصی‌سازی شده در نظر گرفته شود.

به طور هوشمندانه تری عمل کنند. به ویژه، این روش توانایی پردازش داده های ناقص یا نامطمئن را دارد، که در شرایطی که اطلاعات به طور کامل در دسترس نیست، یک مزیت قابل توجه محسوب می شود. در نهایت، این معماری پیش بینی می تواند تجربه آموزشی بهینه ای برای کاربران فراهم کرده و کیفیت خدمات را ارتقا دهد. به این ترتیب، این روش به عنوان یک راهکار جامع در سیستم های آموزش الکترونیکی قابل استفاده است.

کاربران در هر گروه بررسی می شود و بازخوردهای آن ها برای بهبود مدل و به روزرسانی شبکه بیزین استفاده می شود. این روش به سیستم این امکان را می دهد که به طور مداوم کیفیت خدمات خود را ارتقا دهد و از بازخوردهای کاربران برای بهبود فرآیندها بهره برداری کند. روش پیشنهادی این پژوهش با استفاده از الگوریتم های یادگیری ماشینی و شبکه های بیزین، امکان شخصی سازی و بهبود تجربه کاربری را فراهم می آورد و به سیستم های آموزش الکترونیکی کمک می کند تا



شکل ۱: معماری روش پیشنهادی

برخورد کمتر و با نزدیک ترین شکست بیشتر باشد، چراکه ویژگی های بهتر به تفکیک کلاس ها کمک می کنند. وزن دهی ویژگی ها بر اساس این اختلاف ها به روز می شود: اختلاف بین نمونه انتخاب شده و نزدیک ترین برخورد، وزن ویژگی را کاهش می دهد، و اختلاف با نزدیک ترین شکست وزن را افزایش می دهد. ویژگی هایی که وزن بالاتری دارند، بهتر می توانند نمونه های مشابه را از هم تفکیک کنند. در پایان، ویژگی هایی که وزن آن ها از یک حد آستانه کمتر است، حذف می شوند و ویژگی های باقی مانده به عنوان مجموعه ویژگی های منتخب استفاده می شوند. این حد آستانه توسط کاربر یا خودکار و با آزمون و خطا تعیین می شود.

۲،۳ انتخاب ویژگی

برای یافتن کاربران مشابه، ابتدا ویژگی های مشترک هر کاربر را انتخاب کرده و سپس از الگوریتم خوشه بندی استفاده می کنیم تا کاربران را در گروه های مشابه قرار دهیم. این ویژگی ها ممکن است بین کاربران متفاوت باشد و به نوع سرویس و اطلاعات موجود برای هر کاربر بستگی دارد. به عنوان روش اصلی برای انتخاب ویژگی، از الگوریتم Relief استفاده می کنیم که مبتنی بر روش های آماری و وزن دهی است. در این روش، ابتدا زیرمجموعه ای از داده های آموزشی انتخاب می شود. سپس برای هر نمونه در این زیرمجموعه، نزدیک ترین نمونه های مشابه (نزدیک ترین برخورد) و متفاوت (نزدیک ترین شکست) بر اساس معیار فاصله اقلیدسی شناسایی می شوند.

نزدیک ترین برخورد: شامل نمونه های هم کلاس است که کمترین فاصله را با نمونه مورد نظر دارند. در مقابل، نزدیک ترین شکست: شامل نمونه های خارج از کلاس نمونه است که کمترین فاصله را دارند. هدف از این کار، پیدا کردن ویژگی هایی است که اختلاف آن ها با نزدیک ترین

```

Relief(D, S, NoSample, Threshold)
(1)  $T = \phi$ 
(2) Initialize all weights,  $W_i$ , to zero.
(3) For  $i = 1$  to  $NoSample/*$  Arbitrarily chosen  $*/$ 
    Randomly choose an instance  $x$  in  $D$ 
    Finds its  $nearHit$  and  $nearMiss$ 
    For  $j = 1$  to  $N$ 
         $W_j = W_j - diff(x_j, nearHit_j)^2 + diff(x_j, nearMiss_j)^2$ 
(4) For  $j = 1$  to  $N$ 
    If  $W_j \geq Threshold$ 
        Append feature  $f_j$  to  $T$ 
(5) Return  $T$ 
    
```

شکل ۲: الگوریتم Relief جهت انتخاب ویژگی و کاهش ابعاد

الگوریتم Relief مطابق شکل (۲) برای ویژگی‌های دارای نویز یا ویژگی‌های دارای همبستگی خوب کار می‌کند و پیچیدگی زمانی آن به صورت خطی و تابعی از تعداد ویژگی‌ها و تعداد نمونه‌های مجموعه نمونه می‌باشد، و هم برای داده‌های پیوسته و هم برای داده‌های صوری خوب کار می‌کند. برای کاربران ویژگی‌های زیر در نظر گرفته شده است، برای این که عملیات مربوط به پردازش راحت‌تر و بهتر صورت بگیرد هر یک از ویژگی‌ها به یک رنج عددی بر اساس نوع آن نگاشت شده‌اند. به عنوان مثال جدول (۱) با تمام ویژگی‌ها برای ۱۰ کاربر وجود دارد که ویژگی‌ها با F_1 تا F_{18} نشان داده شده‌اند. در جدول (۱) تمام اعداد فرضی هستند و هدف فقط تشریح روش پیشنهادی برای انتخاب ویژگی می‌باشد و فرض شده است که هر کاربر حداکثر ۱۸ ویژگی دارد.

الگوریتم Relief در این مقاله به منظور کاهش ابعاد داده‌ها با انتخاب ویژگی‌های مهم مورد استفاده قرار گرفته است. فرایند این الگوریتم برای داده‌های جدول (۱) به صورت زیر انجام می‌شود:

جدول ۱: ویژگی‌های برای کاربران مختلف

کد کاربر	کد سرویس	F_1	F_2	F_3	F_4	F_5	F_6	F_7	F_8	F_9	F_{10}	F_{11}	F_{12}	F_{13}	F_{14}	F_{15}	F_{16}	F_{17}	F_{18}
B ₁	A	۵	۱۲	۱۰۰	-	۶۵	۷۹	-	-	۲۵	-	-	۱۲	-	۳۶	-	-	۱۲	۲۰
B ₂	B	۱۰	۱۱	-	۲۵	۷۰	-	۴۵	-	۱۹	۳۵	-	-	۱۴	-	۳	۲	۲۹	-
B ₃	A	۵	-	۱۱۰	۲۹	-	۸۰	۴۹	۸۹	-	۳۶	۷۴	-	۱۵	-	-	-	۱۹	۲۵
B ₄	C	-	۱۹	-	-	۵۵	۸۱	-	-	-	۲۰	-	۳۶	۱۵	-	۳۷	۷	-	-
B ₅	C	۹	۱۵	۱۱۵	۳۰	۷۵	-	-	-	-	-	-	۳۹	-	-	-	۴	۲۰	۲۱
B ₆	A	۸	-	-	۲۸	-	۸۰	۴۵	۹۰	۲۰	-	۷۴	۱۷	-	۳۶	-	-	۲۱	-
B ₇	B	۴	۱۳	۱۲۰	-	۷۵	۸۲	۴۵	۹۱	-	-	۷۲	۱۵	-	۳۶	۷	۷	-	۲۵
B ₈	B	-	۱۴	۱۱۲	-	۷۰	۷۷	۴۵	۱۹	-	۳۶	-	-	۱۵	۳۵	-	۷	-	-
B ₉	C	۲	۱۳	-	۳۰	۷۹	-	-	-	۲۰	-	۷۵	-	۱۵	۳۵	۴	۵	-	۲۲
B ₁₀	A	۸	-	۱۱۳	-	-	۸۰	-	۹۰	۲۰	۳۵	-	۱۸	۱۵	۳۳	۵	-	۲۰	۲۰

• کاربر B_۶ دارای ۱۰ ویژگی می‌باشد.

• کاربر B_{۱۰} دارای ۱۱ ویژگی می‌باشد.

پس سرویس A برای تمام ۴ کاربر این سرویس در مجموع ۴۳ ویژگی و به طور متوسط ۱۰ ویژگی دارد، طبق توضیحات بالا برای تمام سرویس‌ها در جدول (۲) قابل مشاهده می‌باشد و تعداد ویژگی‌های هر سرویس برابر با تعداد میانگین سرویس‌ها به ازای کاربران سیستم خواهد بود.

نوع انتخاب ویژگی بستگی به کاربر دارد زیرا اطلاعات موجود مربوط به انتخاب سرویس‌ها می‌باشد در ابتدا میانگین تعداد ویژگی‌های هر سرویس محاسبه می‌شود در واقع بررسی می‌شود که به طور میانگین برای هر سرویس برای تمام کاربرها چند ویژگی وجود دارد مثلاً برای سرویس A تعداد ویژگی‌های دارای مقدار برای هر کاربر به صورت زیر می‌باشد:

• کاربر B_۱ دارای ۱۰ ویژگی می‌باشد.

• کاربر B_۳ دارای ۱۱ ویژگی می‌باشد.

جدول ۲: میانگین تعداد ویژگی برای هر سرویس

سرویس	میانگین تعداد ویژگی
A	۱۰
B	۱۱
C	۱۰

در ادامه باید ویژگی های انتخاب شده قطعی مشخص شود، این ویژگی ها، ویژگی هایی هستند که در تمام کاربران سرویس مورد نظر دارای مقدار هستند، به عنوان مثال در جدول (۳) کاربرانی که سرویس A را دارند قابل مشاهده اند و ستون های سبز رنگ ویژگی های قطعی هستند زیرا در تمام کاربران دارای ویژگی می باشند.

جدول ۳: ویژگی های قطعی

کد کاربر	کد سرویس	F _۱	F _۲	F _۳	F _۴	F _۵	F _۶	F _۷	F _۸	F _۹	F _{۱۰}	F _{۱۱}	F _{۱۲}	F _{۱۳}	F _{۱۴}	F _{۱۵}	F _{۱۶}	F _{۱۷}	F _{۱۸}
B _۱	A	۵	۱۲	۱۰۰	-	۶۵	۷۹	-	-	۲۵	-	-	۱۲	-	۳۶	-	-	۱۲	۲۰
B _۳	A	۵	-	۱۱۰	۲۹	-	۸۰	۴۹	۸۹	-	۳۶	۷۴	-	۱۵	-	-	-	۱۹	۲۵
B _۶	A	۸	-	-	۲۸	-	۸۰	۴۵	۹۰	۲۰	-	۷۴	۱۷	-	۳۶	-	-	۲۱	-
B _{۱۰}	A	۸	-	۱۱۳	-	-	۸۰	-	۹۰	۲۰	۳۵	-	۱۸	۱۵	۳۳	۵	-	۲۰	۲۰

جدول ۴: ویژگی های انتخاب شده قطعی

سرویس	ویژگی
A	F _۱ -F _۶ -F _{۱۷}
B	F _۲ -F _۵ -F _۷ -F _{۱۶}
C	F _۲ -F _۵ -F _{۱۵}

تا به اینجا ویژگی های قطعی برای هر سرویس انتخاب شده است و در ادامه ویژگی هایی، برای هر سرویس انتخاب خواهند شد که حداقل نصف کاربران دارای ویژگی فوق باشند. در این انتخاب برای هر ویژگی بررسی می شود که چه تعداد از کاربران سرویس مورد نظر را دارند، ویژگی

که ۵۰ درصد کاربران آن را داشته باشند انتخاب خواهند شد به عنوان مثال در جدول (۵) برای سرویس A ویژگی های به رنگ آبی انتخاب می شوند زیرا ۵۰ درصد کاربران حداقل ویژگی فوق را دارند.

جدول ۵: کاربران دارای ویژگی

کد کاربر	کد سرویس	F _۱	F _۲	F _۳	F _۴	F _۵	F _۶	F _۷	F _۸	F _۹	F _{۱۰}	F _{۱۱}	F _{۱۲}	F _{۱۳}	F _{۱۴}	F _{۱۵}	F _{۱۶}	F _{۱۷}	F _{۱۸}
B _۱	A	۵	۱۲	۱۰۰	-	۶۵	۷۹	-	-	۲۵	-	-	۱۲	-	۳۶	-	-	۱۲	۲۰
B _۳	A	۵	-	۱۱۰	۲۹	-	۸۰	۴۹	۸۹	-	۳۶	۷۴	-	۱۵	-	-	-	۱۹	۲۵
B _۶	A	۸	-	-	۲۸	-	۸۰	۴۵	۹۰	۲۰	-	۷۴	۱۷	-	۳۶	-	-	۲۱	-
B _{۱۰}	A	۸	-	۱۱۳	-	-	۸۰	-	۹۰	۲۰	۳۵	-	۱۸	۱۵	۳۳	۵	-	۲۰	۲۰

به عنوان مثال در جدول (۵) ویژگی F_۳ انتخاب می شود زیر ۳ کاربر از ۴ کاربر ویژگی فوق را دارد ولی ویژگی F_{۱۶} انتخاب نمی شود، در جدول (۶) ویژگی های انتخاب شده برای هر سرویس مشخص شده است.

۳.۳ خوشه بندی

برای هر سرویس به صورت مجزا خوشه بندی انجام می شود که ویژگی های برابر دارند. الگوریتم خوشه بندی K-means. داده ها را به K خوشه

جدول ۶: ویژگی های انتخاب شده برای هر سرویس

کد سرویس	F _۱	F _۲	F _۳	F _۴	F _۵	F _۶	F _۷	F _۸	F _۹	F _{۱۰}	F _{۱۱}	F _{۱۲}	F _{۱۳}	F _{۱۴}	F _{۱۵}	F _{۱۶}	F _{۱۷}	F _{۱۸}
A	√	√	-	√	√	√	√	√	√	√	√	√	√	√	-	-	√	√
B	√	√	√	-	√	√	√	-	√	√	-	-	√	√	√	√	-	-
C	√	√	-	√	√	-	-	-	-	√	√	-	√	√	√	√	-	√

برای محاسبه فاصله برای تمام بردارها از بردار میانگین طبق رابطه (۲) استفاده می کنیم.

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2}$$

(۲)

هر چقدر مقدار فاصله کمتر باشد بردارها مشابه تر می باشند. برای مثال فاصله کاربر B_1 با مشخصات زیر:

کد کاربر	F _۱	F _۲	F _۳	F _۴	F _۵
۱B	۱۴	۱۰۶	۳۷	۷۵	۱۵

از بردار میانگین با مشخصات زیر:

کد کاربر	F _۱	F _۲	F _۳	F _۴	F _۵
۰	۷.۱۲	۶۷.۶	۷.۳۶	۱.۷۴	۶.۱۷

به صورت زیر محاسبه می شود.

Distance

$$= \sqrt{\left\{ \begin{aligned} &\left((17.6 - 15)^2 + (74.1 - 75)^2 + (36.7 - 37)^2 \right) \\ &+ \left((6.67 - 6.1)^2 + (12.7 - 4)^2 \right) \end{aligned} \right\}}$$

$$= 3.11045$$

نتیجه محاسبه فاصله برای تمام داده ها از بردار میانگین در جدول (۷) نشان داده شده است.

جدول ۷: میزان فاصله

کد کاربر	فاصله
B_1	۱۱۰۴۵.۳
B_2	۲۳۷۴۹.۲۱
B_3	۱۸۸۸.۱۲
B_4	۰۲۰۱۲.۱۰
B_5	۷۵۷۱۳۷.۶
B_6	۹۵۳۵۳.۱۱
B_7	۱۸۸۸.۱۲
B_8	۱۴۲۰۹.۱۵
B_9	۷۵۷۱۳۷.۶
B_{10}	۹۴۰۲۴.۷۴

میانگین فاصله ها برابر است با ۴۲.۱۷ بنابراین داده بردار و B_8 و B_{10} به عنوان سرخوشه انتخاب می شوند زیرا B_{10} نسبت به مقدار میانگین ۴۲.۱۷ بیشترین فاصله مثبت را دارد.

B_1 نسبت به مقدار میانگین ۴۲.۱۷ بیشترین فاصله منفی را دارد (کوچکترین عدد در اعداد کوچکتر از ۴۲.۱۷). B_8 کمترین فاصله را از بردار میانگین دارد. برای انتخاب اعضای هر خوشه طبق رابطه گفته شده، فاصله هر داده از سرخوشه های تعیین شده محاسبه می شود. بردارها جذب خوشه های می شوند که کمترین فاصله را دارند که نتیجه در جدول (۸) قابل مشاهده است.

جدول ۸: نتیجه خوشه بندی

اعضا	سرخوشه
B_5, B_9	B_1
B_7, B_2, B_4, B_6	B_8
B_3	B_{10}

برای داده های بعدی که وارد می شوند فاصله هر بردار از سرخوشه ها محاسبه می شود و اگر فاصله بردار از میانگین فاصله ها کمتر باشد جذب یکی از سرخوشه ها می شود و در غیر این صورت خود به عنوان یک سرخوشه جدید انتخاب می شود.

۴. پیشنهاد سرویس به کاربر

۱.۴ انتخاب سرویس کاندید

روش پیشنهادی برای پیشنهاد سرویس به کاربر شامل دو حالت است: وضعیت اول: اگر کاربر یک سرویس را انتخاب کند، سیستم بر اساس آن سرویس، سایر سرویس های مرتبط را پیشنهاد می دهد. این وضعیت شامل سه مرحله است:

- ایجاد جدولی از پایگاه داده که سرویس های مرتبط با هدف کاربر را نشان می دهد. سپس کاربرانی که سرویس مورد نظر را انتخاب کرده اند، شناسایی و سایر کاربران حذف می شوند. حد آستانه برای انتخاب سرویس ها، نصف تعداد این کاربران در نظر گرفته می شود.

- سرویس هایی که فراوانی آن ها از حد آستانه بیشتر است انتخاب می شوند. سپس مجموعه های بزرگتری از این سرویس ها ساخته شده و مجموعه هایی که از آستانه کمتر باشند حذف می شوند.

- مجموعه های باقی مانده براساس معیارهای اطمینان و پوشش بررسی می شوند تا سرویس های پیشنهادی نهایی انتخاب شوند.

وضعیت دوم: اگر کاربر هیچ سرویسی انتخاب نکرده باشد، سیستم با توجه به تشابه او با دیگر کاربران، سرویس های مناسب را پیشنهاد می دهد. این وضعیت جهت پیشنهاد سرویس به کاربر شامل مراحل زیر می باشد:

- کاربر عنوان هدف را انتخاب می کند.
 - سیستم جدول مربوط به هدف مورد نظر را می سازد.
 - تمام قوانین ممکن را با توجه به حد آستانه تولید می کند.
 - قوانینی که از نظر پوشش و اطمینان قابل قبول باشند به عنوان پیشنهاد به کاربر ارائه می شوند.
- برای فهم بیشتر مطلب در اینجا یک نمونه پیشنهاد انتخاب سرویس به کاربر را مورد بررسی قرار می دهیم. در جدول (۹) نشان می دهد تعداد ۶ سرویس برای هدف مورد نظر کاربر وجود دارد و چه کاربرهای چه سرویس هایی را انتخاب کرده اند.

جدول ۹: سرویس های انتخابی برای هدف مورد نظر

کاربر	A	B	C	D	E	F
B_1	*	*		*	*	*
B_2		*	*		*	*
B_3	*		*	*		*
B_4		*	*	*	*	
B_5	*			*		*

بررسی می شود که X کاربران می باشد و A سرویس مورد نظر می باشد. در این مثال درجه اطمینان ۵۰٪ در نظر گرفته شده است. رابطه پوشش و اطمینان به صورت زیر تعریف می شوند.

تعریف پوشش^۱: مجموعه I شامل تمام آیتم ها و مجموعه آیتم X را در نظر بگیرید، می گوئیم پوشش X در پایگاه داده برابر است با l اگر و فقط اگر تعداد وقوع مجموعه آیتم X در پایگاه داده برابر با l باشد.

درجه اطمینان^۲: مجموعه I شامل تمامی اقلام و مجموعه آیتم های X و

Y مفروض اند. درجه اطمینان قانون $X \Rightarrow Y$ برابر است با: $X \cap Y = \emptyset$:

$$S \text{ Conf}(X \rightarrow Y) = \text{support}(X \cup Y) / \text{Support}(X)$$

$$\text{Support}(X \rightarrow Y) = \text{support}(X \cup Y) / m$$

در رابطه بالا m تعداد کل رکوردهای Answer می باشد.

جدول ۱۱: قوانین اطمینان و پوشش

مجموعه	قانون	اطمینان	پوشش	وضعیت
{A,B}	$\{A \rightarrow B\}$	۳۶.۰	۲۳.۰	ضعیف
{A,C}	$\{A \rightarrow C\}$	۱۸.۰	۱۲.۰	ضعیف
{A,D}	$\{A \rightarrow D\}$	۵۴.۰	۳۵.۰	قوی
{A,E}	$\{A \rightarrow E\}$	۳۶.۰	۲۳.۰	ضعیف
{A,F}	$\{A \rightarrow F\}$	۱۸.۰	۱۲.۰	ضعیف
{A,B,D}	$\{A \rightarrow B,D\}$	۱۸.۰	۱۲.۰	ضعیف
{A,B,E}	$\{A \rightarrow B,E\}$	۱۸.۰	۱۲.۰	ضعیف
{A,D,E}	$\{A \rightarrow D,E\}$	۱۸.۰	۱۲.۰	ضعیف
{A,C,D}	$\{A \rightarrow C,D\}$	۰۹.۰	۰۶.۰	ضعیف
{A,D,F}	$\{A \rightarrow D,F\}$	۰۹.۰	۰۶.۰	ضعیف
{A,B,D,E}	$\{A \rightarrow B,D,E\}$	۰۹.۰	۰۶.۰	ضعیف

با توجه به نتایج خروجی در جدول (۱۰) سرویس های آموزشی که در دو معیار $\sup$ و conf مقدار بالای ۰/۵ دارند به عنوان سرویس های مطلوب و کاندید انتخاب می شوند و در ادامه با استفاده از الگوریتم بیزین بررسی می شوند که آیا برای کاربر مورد نظر دارای کیفیت هستند یا خیر و بهترین سرویس را به کاربر پیشنهاد می دهد.

۲.۴ انتخاب بهترین سرویس کاندید مبتنی بر شبکه بیزین

تا به اینجا با استفاده از الگوریتم گفته شده سرویس های کاندید برای پیشنهاد به کاربر آموزشی انتخاب شده اند و در این قسمت با استفاده از الگوریتم بیزین بهترین سرویس ها به کاربر ارائه خواهد شد. در این قسمت ورودی الگوریتم بیزین سرویس های آموزشی کاندید هستند و خروجی آن تعیین این موضوع که آیا این سرویس ها برای کاربر هدف از کیفیت قابل قبولی برخوردار هستند یا خیر.

الگوریتم بیزین برای انجام این کار از مجموعه داده آموزش استفاده می کند، در این مجموعه داده تعدادی سرویس آموزش با ویژگی های مختلف وجود دارد که مشخص شده است که دارای کیفیت هستند یا خیر، سرویس های با کیفیت دارای کلاس یک و سرویس های بدون کیفیت دارای کلاس صفر هستند که در جدول زیر قابل مشاهده می باشد.

B_6	*	*	*	*	*	*
-------	---	---	---	---	---	---

حال فرض کنید چهار کاربر سرویس A را انتخاب کرده اند و حد آستانه ۲ در نظر گرفته می شود. مطابق جدول (۱۰) فقط انتخاب کاربرانی که سرویس A را انتخاب کرده اند، بررسی می شود.

جدول ۱۰: انتخاب سرویس کاربرانی که حتماً A را انتخاب کرده اند

کاربر	A	B	C	D	E	F
B_1	*	*		*	*	*
B_3	*		*	*		*
B_5	*			*		*
B_6	*	*	*	*	*	*

ابتدا فراوانی تک تک سرویس ها محاسبه می شود.

$$\{A\}=4, \{B\}=2, \{C\}=2, \{D\}=4, \{E\}=2, \{F\}=3$$

به دلیل اینکه فراوانی همه سرویس ها از حد آستانه بالاتر هستند

سرویس حذف نمی شود و مجموعه L_1 ساخته می شود.

$$L_1 = \{\{A\}, \{B\}, \{C\}, \{D\}, \{E\}, \{F\}\}$$

مجموعه دو عضوی با استفاده از L_1 ساخته می شود.

- $\{A,B\}=2 \quad \{A,C\}=2 \quad \{A,D\}=4 \quad \{A,E\}=2 \quad \{A,F\}=3$
- $\{B,C\}=1 \quad \{B,D\}=2 \quad \{B,E\}=2 \quad \{B,F\}=1$
- $\{C,D\}=2 \quad \{C,E\}=1 \quad \{C,F\}=1$
- $\{D,E\}=2 \quad \{D,F\}=3$
- $\{E,F\}=1$
- $L_2 = \{\{A,B\}, \{A,C\}, \{A,D\}, \{A,E\}, \{A,F\}, \{B,D\}, \{B,E\}, \{C,D\}, \{D,E\}, \{D,F\}\}$

با استفاده از مجموعه L_2 مجموعه L_3 ساخته می شود.

- $\{A,B,C\}=1 \quad \{A,B,D\}=2 \quad \{A,B,E\}=2 \quad \{A,B,F\}=1$
- $\{A,C,D\}=2$
- $\{A,C,E\}=1 \quad \{A,C,F\}=1 \quad \{A,D,E\}=2 \quad \{A,D,F\}=3$
- $\{A,E,F\}=1$
- $\{B,D,E\}=2 \quad \{B,D,C\}=1 \quad \{B,D,E\}=2 \quad \{B,D,F\}=1$
- $L_3 = \{\{A,B,D\}, \{A,B,E\}, \{A,C,D\}, \{A,D,E\}, \{A,D,F\}, \{B,D,E\}\}$

با استفاده از مجموعه L_3 مجموعه L_4 ساخته می شود:

- $\{A,B,D,E\}=2 \quad \{A,B,C,D\}=1 \quad \{A,B,D,F\}=1 \quad \{A,C,D,E\}=1$
- $\{A,C,D,F\}=1 \quad \{A,C,D,E\}=1 \quad \{A,D,E,F\}=1$
- $L_4 = \{\{A,B,D,E\}\}$

مجموع نهایی از L_4, L_3, L_2 به صورت زیر می باشد:

$$\text{Answer} = \{\{A,B\}, \{A,C\}, \{A,D\}, \{A,E\}, \{A,F\}, \{B,D\}, \{B,E\}, \{C,D\}, \{D,E\}, \{D,F\}, \{A,B,D\}, \{A,B,E\}, \{A,C,D\}, \{A,D,E\}, \{A,D,F\}, \{B,D,E\}, \{A,B,D,E\}\}$$

سپس باید میزان اطمینان و پوشش را برای مجموعه فوق محاسبه کرد. در این محاسبه قانون $A \rightarrow X$ باید از نظر اطمینان و پوشش

² - confidence

¹ - support

در جدول (۱۲) تعداد ۱۵ بردار آموزش به عنوان مثال نمایش داده شده- است. در این جدول هر ردیف یک سرویس آموزش می باشد که دارای ویژگی های F_1 تا F_5 می باشد.

جدول ۱۲: داده های آموزش

class	F _۱	F _۲	F _۳	F _۴	F _۵
۱	۵	۷	۹	۶	۵
۰	۵	۲	۳	۴	۷
۰	۶	۲	۵	۴	۳
۱	۱	۱	۲	۴	۶
۰	۳	۲	۴	۷	۷
۱	۵	۳	۳	۶	۵
۰	۴	۴	۲	۱	۴
۱	۲	۳	۵	۴	۲
۰	۴	۱	۴	۶	۵
۱	۲	۳	۲	۴	۲
۱	۳	۲	۱	۳	۱
۱	۲	۳	۵	۵	۱
۰	۱	۵	۳	۷	۴
۱	۲	۶	۲	۶	۳
۰	۵	۷	۱	۴	۲

در این جا هدف ما این است با الگوریتم بیزین مشخص کنیم که سرویس های جدول تست دارای کیفیت هستند یا خیر؟ شبکه بیزین یک فرآیند یادگیری بر اساس نظریه یادگیری آماری است که یکی از بهترین روش های یادگیری ماشین مورد استفاده در داده کاوی می باشد. در مسائل مربوط به طبقه بندی داده ها و شناسایی الگو، همچون دسته بندی متون، تشخیص چهره در تصاویر، تشخیص ارقام دست نویس و بیوانفورماتیک کارکرد بسیار موفق داشته است. شبکه بیزین در واقع یک طبقه بندی کننده دودویی است که دو کلاس را با استفاده از یک مرز خطی از هم جدای می کند. این روش با استفاده از یک الگوریتم بهینه سازی، نمونه هایی که مرزهای کلاس ها را تشکیل می دهند به دست می آورند. این نمونه ها را بردارهای پشتیبان گویند. تعدادی از نقاط آموزشی که کمترین فاصله تا مرز تصمیم گیری را دارند می توانند به عنوان زیر مجموعه ای برای تعریف مرزهای تصمیم گیری و به عنوان بردار پشتیبان در نظر گرفته شوند. فرض کنید داده ها از دو کلاس تشکیل شده و کلاس ها در مجموع دارای $L \times X_i = i, i = 1, \dots$ نقطه آموزشی باشند که X_i یک بردار است. برای محاسبه مرز تصمیم گیری دو کلاس کاملاً جدا از هم از روش حاشیه بهینه استفاده می شود. در این روش مرز خطی بین دو کلاس به گونه ای محاسبه می شود که ابتدا: تمام نمونه های کلاس مثبت ۱ در یک طرف مرز و تمام نمونه های کلاس منفی ۱ در طرف دیگر مرز واقع شوند. سپس مرز تصمیم گیری به گونه ای باشد که فاصله نزدیک ترین نمونه های آموزشی هر دو کلاس از یکدیگر در راستای عمود بر مرز تصمیم گیری تا جایی که ممکن است حداکثر شود. یک مرز تصمیم گیری خطی را در حالت کلی می توان به صورت زیر نوشت:

$$w \cdot x + b =$$

X یک نقطه روی مرز تصمیم گیری و w یک بردار n بعدی عمود بر مرز تصمیم گیری است. $b/||w||$ فاصله مبدأ تا مرز تصمیم گیری و w بیانگر ضرب داخلی دو بردار w و x است. از آنجا که با ضرب یک ثابت در دو طرف معادله باز هم تساوی برقرار خواهد بود. اولین مرحله برای محاسبه مرز تصمیم گیری بهینه، پیدا کردن نزدیک ترین نمونه های آموزشی دو کلاس است. در مرحله بعد فاصله آن نقاط از هم در راستای عمود بر مرزهایی که دو کلاس را به طور کامل جدای می کنند محاسبه می شود. مرز تصمیم گیری بهینه مرزی است که حداکثر حاشیه را داشته باشد.

۵. آزمایش ها

۵.۱. آزمایش اول

در این آزمایش رده بند بر اساس داده های غیر نرمال قسمت آموزش عمل می کند. پارامتر اصلی که روی کارایی رده بندی تأثیر زیادی دارد و عنصری مهم در قابلیت یادگیری است، شعاع نمونه های خودی یا RS^2 است. درصد تشابه هر داده تست با داده های قسمت آموزش که در وضعیت غیر نرمال دارند محاسبه می شوند. درصد تشابه برابر است با تعداد ویژگی های مشترک که بردار با بردارهای داده آموزش داشته باشد؛ که در آخر به صورت درصد نمایش داده می شود و بالاترین مقدار به عنوان درصد تشابه انتخاب می شود. مقدار RS برای یک بردار برابر با میانگین مقدار درصد تشابه کل بردارها است.

خروجی این رده بند زمانی یک می شود که درصد تشابه یک بردار از مقدار RS بیشتر باشد. در جدول (۱۳) مقدار درصد تشابه برای داده های تست با استفاده از داده های غیر نرمال آمده است.

جدول ۱۳: مقدار درصد تشابه داده های تست با داده های غیر نرمال

درصد تشابه با داده های غیر نرمال	F _۱	F _۲	F _۳	F _۴	F _۵
%۴۰	۳	۱	۶	۷	۵
%۶۰	۴	۲	۳	۴	۵
%۲۰	۲	۵	۲	۲	۱
%۴۰	۶	۴	۴	۲	۳
%۴۰	۵	۳	۵	۱	۴
%۴۰	۷	۱	۲	۵	۴

با توجه به درصد تشابه در جدول (۱۳) برای داده تست که بر اساس داده های غیر نرمال داده آموزش به دست آمده اند مقدار RS برابر است با:

$$RS = (0.4 + 0.6 + 0.2 + 0.4 + 0.4 + 0.4) / 6 \rightarrow RS = 0.4\%$$

بنا بر آنچه گفته شد بردارهایی که مقدار تشابه آن ها از RS بیشتر باشد خروجی رده بند ۱ برای آن ها یک می شود.

۵.۲. آزمایش دوم

که در ادامه توضیح داده می شود، استفاده می کنیم. در جدول (۱۵) نتیجه ارزیابی اولیه رده بندها نمایش داده شده است.

جدول ۱۵: نتیجه ارزیابی اولیه

class	خروجی رده بند ۲	خروجی رده بند ۱	F۱	F۲	F۳	F۴	F۵
غیر نرمال	۰	۱	۳	۱	۶	۷	۵
همپوشانی	۱	۱	۴	۲	۳	۴	۵
نرمال	۱	۰	۲	۵	۲	۲	۱
غیر نرمال	۰	۱	۶	۴	۴	۲	۳
همپوشانی	۱	۱	۵	۳	۵	۱	۴
همپوشانی	۱	۱	۷	۱	۲	۵	۴

بردارهای که class آن ها نرمال یا غیرنرمال شده است وضعیت مشخص دارند ولی برای مواردی که مقدار همپوشانی دارند طبق روش های گفته شده در زیر عمل می کنیم.

۳.۵ حالت همپوشانی (مبهم)

در آزمایش اول و دوم این حالت زمانی رخ می دهد که خروجی هر دو رده بند ۱ یا ۰ باشد و در این وضعیت از آمارگیری استفاده می کنیم. تعداد داده بالغ گروهی که شباهت بیشتری به داده تست داشته باشند برچسب آن برای داده تست انتخاب می شود؛ و اگر درصد تشابه برای هر دو گروه نرمال و غیرنرمال برابر بود به صورت زیر عمل می کند. در این روش هر ویژگی بردار اگر تنها یک اختلاف با ویژگی بردار ناحیه آموزش داشته باشد مشابه در نظر گرفته می شود. در جدول (۱۶) چند بردار در حالت همپوشانی قرار گرفته اند که برای مشخص کردن وضعیت نرمال و غیرنرمال آن ها از روش گفته شده استفاده می کنیم که به صورت بردار به بردار در ادامه توضیح داده شده است. بردارهای که در وضعیت همپوشانی قرار دارند با درصد تشابه آن ها نمایش داده شده است.

جدول ۱۶: وضعیت بردارهای که حالت همپوشانی دارند

class	درصد تشابه با داده های غیر نرمال	درصد تشابه با داده های نرمال	F۱	F۲	F۳	F۴	F۵
غیر نرمال	۸۰٪	۶۰٪	۴	۲	۳	۴	۵
غیر نرمال	۸۰٪	۶۰٪	۵	۳	۵	۱	۴
غیر نرمال	۸۰٪	۶۰٪	۷	۱	۲	۵	۴

حالت دوم: در این حالت کاربر سرویسی انتخاب نمی کند و فقط هدف مورد نظر را مشخص می کند و سیستم چند سرویس با توجه به اطلاعات قبلی به کاربر پیشنهاد می دهد. اگر کاربر سرویسی را انتخاب کند که قبلاً هیچ کاربری دیگری آن را انتخاب نکرده باشد سیستم با حالت دوم خود سرویس به او پیشنهاد می دهد. روش کار به ترتیب زیر می باشد:

- کاربر عنوان هدف را انتخاب می کند.

در این آزمایش رده بند بررسی می کند که آیا بردارهای فعلی سرویس آموزشی در وضعیت نرمال قرار دارند یا نه. این رده بند بر اساس داده های نرمال قسمت آموزش عمل می کند و اگر بردار نرمال باشد خروجی یک و در غیر این صورت خروجی صفر می شود. همانند رده بند اول باید مقدار RS برای این قسمت نیز محاسبه شود. محاسبه RS برای کل داده های تست بدین صورت می باشد.

درصد تشابه هر داده تست با داده های قسمت آموزش که در وضعیت نرمال دارند محاسبه می شوند. درصد تشابه برابر است با تعداد ویژگی های مشترک که بردار با بردارهای داده آموزش داشته باشد؛ که در آخر به صورت درصد نمایش داده می شود و بالاترین مقدار به عنوان درصد تشابه انتخاب می شود. مقدار RS برای یک بردار برابر با میانگین مقدار درصد تشابه کل بردارها است. خروجی این رده بند زمانی یک می شود که درصد تشابه یک بردار از مقدار RS بیشتر باشد. در جدول (۱۴) درصد تشابه داده های تست با داده های نرمال قسمت آموزش نمایش داده شده است.

جدول ۱۴: مقدار درصد تشابه داده های تست با داده های نرمال

درصد تشابه با داده های نرمال	F۱	F۲	F۳	F۴	F۵
۲۰٪	۳	۱	۶	۷	۵
۴۰٪	۴	۲	۳	۴	۵
۴۰٪	۲	۵	۲	۲	۱
۲۰٪	۶	۴	۴	۲	۳
۴۰٪	۵	۳	۵	۱	۴
۴۰٪	۷	۱	۲	۵	۴

با توجه به درصد تشابه در جدول (۱۴) برای داده تست که بر اساس داده های نرمال داده آموزش به دست آمده اند مقدار RS برابر است با:

$$RS = \frac{(0.2 + 0.4 + 0.4 + 0.2 + 0.4 + 0.4)}{6} \rightarrow RS = 0.33$$

بنا بر آنچه گفته شد بردارهایی که مقدار تشابه آن ها از RS بیشتر باشد خروجی رده بند ۲ برای آن ها یک می شود. مکانیسم انتخاب بهترین رده بند در فاز تست برای یک نمونه جدید ورودی، چهار حالت رخ می دهد:

در آزمایش اول اگر رده بند که ناحیه غیر خودی را می پوشاند خروجی یک داشته باشد؛ و در آزمایش دوم رده بند که ناحیه خودی را پوشش می دهد، خروجی صفر داشته باشد. در این حالت با قطعیت ۱۰۰٪ می توان گفت نمونه جدید، متعلق به کلاس دو، یعنی یک نمونه ناهنجاری است و سیستم پیغام هشدار تولید می کند. همچنین اگر در آزمایش اول رده بند خروجی صفر داشته باشد و در آزمایش دوم رده بند خروجی یک داشته باشد. در این حالت با قطعیت ۱۰۰٪ می توان گفت نمونه جدید، متعلق به کلاس یک یعنی یک نمونه نرمال است.

و نهایتاً اگر هر دو آزمایش اول و دوم رده بند خروجی صفر داشته باشند یا هر دو رده بند خروجی یک داشته باشند، اصطلاحاً چنین حالتی را همپوشانی می گوئیم و برای رفع آن از روش انتخاب یکی از رده بندها

- سیستم جدول مربوط به هدف مورد نظر را می‌سازد.
- تمام قوانین ممکن را با توجه به حد آستانه تولید می‌کند.
- قوانینی که از نظر پوشش و اطمینان قابل قبول باشند به عنوان پیشنهاد به کاربر ارائه می‌شوند.

۶. ارزیابی

در این پژوهش برای شبیه‌سازی از زبان متلب استفاده شده است و برای هر معیار ارزیابی چندین مرحله در نظر گرفته شده است و در هر مرحله اندازه مجموعه داده بزرگ‌تر می‌شود. برای استخراج نتایج هر مرحله چندین مرتبه شبیه‌سازی اجرا شده است و میانگین نتایج به عنوان نتیجه خروجی در نظر گرفته شده است. برای ارزیابی کارایی سیستم‌های طبقه‌بندی معیارهای سنجش متفاوتی وجود دارد. از جمله این معیارها می‌توان به دقت طبقه‌بندی، فراخوانی، صحت و نرخ خطا اشاره کرد. ابتدا پارامترهایی که در تعریف این معیارها استفاده می‌شوند، معرفی می‌کنیم: مثبت حقیقی (TP): برابر است با درصدی از تعداد اعضای از یک کلاسی مانند X که سیستم طبقه‌بندی کننده آن‌ها را به درستی عضو کلاس X طبقه‌بندی کرده است.

منفی حقیقی (TN): برابر است با درصدی از تعداد اعضای کلاسهای دیگر که عدم تعلق آن‌ها به کلاس X به درستی طبقه‌بندی شده است. مثبت کاذب (FP): برابر است با درصدی از تعداد اعضای کلاسهای دیگر که به اشتباه به عنوان عضو از کلاس X طبقه‌بندی شده است. منفی کاذب (FN): برابر است با درصدی از تعداد اعضای کلاس X که به اشتباه به عنوان عضو از کلاسهای دیگر طبقه‌بندی شده‌اند. مثبت (P): برابر با مجموع تعداد اعضای از کلاس‌ها است که به صورت درست طبقه‌بندی شده‌اند. منفی (N): برابر با مجموع تعداد اعضای از کلاس‌ها است که به صورت اشتباه طبقه‌بندی شده‌اند.

معیار ارزیابی فراخوانی عبارت است از دقت سیستم طبقه‌بندی کننده در طبقه‌بندی درست اعضای از کلاس X که به صورت درست عضو کلاس X طبقه‌بندی شده‌اند و به صورت رابطه ۳ محاسبه می‌شود:

$$\text{Recall} = \frac{TP}{TP + FN} \quad (۳)$$

معیار ارزیابی صحت یعنی اینکه چند درصد از اعضای تشخیص داده شده به عنوان عضو کلاس X در واقع به کلاس X متعلق هستند و به صورت رابطه (۴) محاسبه می‌شود:

$$\text{Precision} = \frac{TP}{TP + FP} \quad (۴)$$

همچنین دقت طبقه‌بندی شاخص دیگری برای ارزیابی عملکرد سیستم‌های طبقه‌بندی کننده می‌باشد و دارای دید و دامنه‌های وسیع‌تر و جامع‌تری نسبت به عملکرد سیستم‌های طبقه‌بندی کننده می‌باشد و عبارت است از نسبت تمام مواردی که به صورت صحیح طبقه‌بندی شده‌اند. دقت طبقه‌بندی به صورت رابطه ۵ محاسبه می‌شود:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (۵)$$

علاوه بر سه روش فوق معیار دیگری به نام F-Score که میانگین وزن دار بین معیارهای صحت و فراخوانی است برای تعیین میزان کارایی سیستم‌های طبقه‌بندی کننده وجود دارد که به صورت رابطه (۶) محاسبه می‌شود.

$$F - \text{score} = ۲ * \frac{\text{Precision} * \text{Recall}}{\text{Precision} + \text{Recall}} \quad (۶)$$

در نهایت میزان خطای سیستم ترکیبی پیشنهادی نیز با استفاده از رابطه (۷) محاسبه می‌گردد:

$$\text{Error rate} = ۱۰۰ - \left(\frac{TP + TN}{TP + TN + FP + FN} \right) \quad (۷)$$

جدول ۱۷: نتایج بهینه‌سازی سیستم‌های آموزشی تحت وب با روش بیزین

F-MESURE	PRECISION	RECALL	FRR-۱	FPR-۱	ACC	زمان اجرا (ثانیه)	تعداد سیستم (سرویس)	تعداد کاربران
۰,۹۱۶۶	۰,۸۶۳۳	۰,۹۷۶۹	۰,۱۲۶۴	۰,۰۱۸۳	۰,۹۴۲۰	۱۷۴	۲۴	۵۰۰
۰,۸۸۸۰	۰,۸۶۰۷	۰,۹۱۷۱	۰,۱۲۷۱	۰,۰۷۱۶	۰,۸۸۹۰	۳۷۹	۲۴	۱۰۰۰
۰,۸۸۲۱	۰,۸۶۵۳	۰,۸۹۹۵	۰,۱۲۲۰	۰,۰۸۷۹	۰,۸۷۲۷	۶۷۴	۲۴	۱۵۰۰
۰,۸۷۹۰	۰,۸۶۱۱	۰,۸۹۷۷	۰,۱۲۶۷	۰,۰۸۹۷	۰,۸۷۱۵	۸۵۸	۲۴	۲۰۰۰

پیچیدگی آن‌ها پیش از ورود به فرآیند خوشه‌بندی و مدل شبکه بیزین است. دقت مدل در تمامی سناریوها بالای ۹۰ درصد گزارش شده که این میزان دقت، تأثیر مثبت کاهش ابعاد داده‌ها در حفظ اطلاعات مهم و حذف نویز را نشان می‌دهد.

مقادیر FRR-1 و FPR-1 که به ترتیب بیانگر کاهش خطاهای مثبت و منفی کاذب هستند، در بازه ۹۵ تا ۹۸ درصد قرار دارند. این مقادیر بالا نشان می‌دهند که سیستم توانسته است هم از پیشنهاد

جدول (۱۷) عملکرد مدل پیشنهادی را در شرایط مختلف از نظر تعداد کاربران و سرویس‌ها ارزیابی می‌کند و نشان می‌دهد که الگوریتم‌های به کاررفته توانسته‌اند تعادلی مناسب میان کارایی، دقت، و زمان اجرا برقرار کنند. زمان اجرا، که نمایانگر بهره‌وری محاسباتی است، با افزایش تعداد کاربران و سرویس‌ها به صورت خطی افزایش یافته، اما همچنان در محدوده‌ای بهینه باقی مانده است. این موضوع نشان‌دهنده کارایی بالای الگوریتم Relief در کاهش ابعاد داده‌ها و مدیریت

۲,۶ معیار دقت

همچنین دقت طبقه‌بندی شاخص دیگری برای ارزیابی عملکرد سیستم‌های طبقه‌بندی کننده می‌باشد و دارای دید و دامنه‌های وسیع‌تر و جامع‌تری نسبت به عملکرد سیستم‌های طبقه‌بندی کننده می‌باشد و عبارت است از نسبت تمام مواردی که به صورت صحیح طبقه‌بندی شده‌اند. دقت طبقه‌بندی به صورت رابطه (۸) محاسبه می‌شود:

$$\text{Recall} = \frac{TP}{TP+FN} \quad (۸)$$

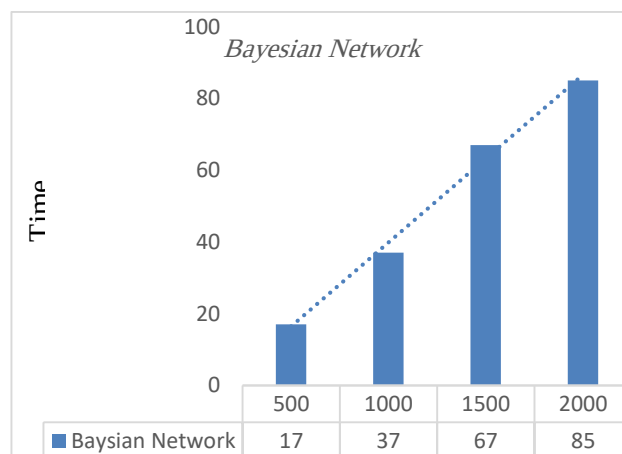
همانطور که در شکل (۴) و (۵) ملاحظه می‌شود نتایج مقایسه با دیگر روش‌ها نشان می‌دهد که دقت بالاتر روش پیشنهادی در مقایسه با روش‌های درخت تصمیم و شبکه‌های عصبی، مستقیماً به ویژگی‌های الگوریتم‌های Relief و K-Means بازمی‌گردد که به‌طور بهینه در فرآیند کاهش ابعاد و خوشه‌بندی کاربران به کار گرفته شده‌اند. الگوریتم Relief با ارزیابی وزن ویژگی‌ها بر اساس تأثیر آن‌ها در تمایز نمونه‌های مشابه و نامشابه، توانسته‌است ویژگی‌های غیرمرتبط و نویزی را از مجموعه داده‌ها حذف کند. این فرآیند، داده‌های ورودی به مدل شبکه بیزین را ساده‌تر و مؤثرتر کرده‌است، به گونه‌ای که تنها ویژگی‌هایی که بیشترین تأثیر را در پیش‌بینی رفتار کاربران داشته‌اند، حفظ شده‌اند. در مقابل، روش‌های درخت تصمیم و شبکه‌های عصبی بدون کاهش مؤثر ابعاد یا تمرکز بر ویژگی‌های کلیدی، از داده‌های خام استفاده کرده‌اند، که این موضوع منجر به کاهش دقت آن‌ها شده‌است. به‌ویژه، در روش درخت تصمیم، وجود نویز یا ویژگی‌های زائد می‌تواند ساختار درخت را پیچیده کرده و منجر به بیش‌برازش شود. همچنین، شبکه‌های عصبی به دلیل وابستگی به تعداد زیادی از ویژگی‌ها، در مواجهه با داده‌های دارای ابعاد بالا و نویز، دچار افت دقت شده‌است. از سوی دیگر، الگوریتم K-Means با خوشه‌بندی کاربران بر اساس رفتارها و ویژگی‌های مشابه، توانسته‌است ساختاری منسجم برای ورودی مدل شبکه بیزین ایجاد کند. این خوشه‌بندی، امکان تفکیک بهتر کاربران و ارائه تحلیل‌های دقیق‌تر را فراهم کرده‌است. علاوه بر این، K-Means با استفاده از فاصله اقلیدسی، کاربران را در خوشه‌هایی گروه‌بندی کرده که در داخل هر خوشه، شباهت‌ها حداکثر و تفاوت‌ها بین خوشه‌ها حداقل است. این ویژگی به مدل شبکه بیزین کمک کرده تا رفتارهای هر خوشه را با دقت بیشتری تحلیل کند و در نتیجه، پیشنهادات مرتبط‌تری برای کاربران ارائه دهد. ترکیب این دو الگوریتم باعث شده که مدل پیشنهادی نسبت به روش‌های سنتی مانند درخت تصمیم و روش‌های پیچیده‌تر مانند شبکه‌های عصبی، از داده‌های تمیزتر و خوشه‌بندی شده استفاده کند. Relief داده‌های ورودی را ساده کرده و K-Means این داده‌های ساده‌شده را در قالب خوشه‌های منطقی مرتب کرده‌است. این فرآیند موجب کاهش خطاهای مثبت و منفی کاذب شده و دقت را در مقایسه با سایر روش‌ها افزایش داده‌است. به‌طور کلی، ترکیب ویژگی‌های Relief و K-Means

سرویس‌های نامرتبط جلوگیری و هم اطمینان حاصل کند که تمامی سرویس‌های مرتبط با کاربران شناسایی می‌شوند. Recall، که درصد سرویس‌های صحیح شناسایی شده را اندازه‌گیری می‌کند، در تمامی سناریوها بالای ۹۵ درصد بوده و نشان می‌دهد که مدل از نظر جامعیت عملکرد، بسیار قوی عمل کرده‌است. از سوی دیگر، Precision، که به کیفیت پیش‌بینی سرویس‌های مرتبط می‌پردازد، در بازه ۹۰ تا ۹۵ درصد باقیمانده و با افزایش تعداد کاربران و سرویس‌ها، کاهش جزئی آن قابل مشاهده است. این کاهش اندک می‌تواند ناشی از افزایش هم‌پوشانی ویژگی‌های کاربران در مقیاس بزرگ‌تر باشد.

شاخص F-MESURE که ترکیبی موزون از Precision و Recall است، در تمامی سناریوها بالای ۹۲ درصد گزارش شده و نشان‌دهنده تعادل مناسب بین جامعیت و دقت مدل است. این مقادیر بالا تأیید می‌کنند که ترکیب الگوریتم‌های Relief و K-Means نه تنها توانسته داده‌های پیچیده را به‌طور مؤثر پردازش کند، بلکه موجب بهبود کارایی مدل شبکه بیزین در پیش‌بینی و ارائه پیشنهادات شخصی‌سازی شده به کاربران شده‌است. به‌طور خاص، افزایش تعداد سرویس‌ها و کاربران منجر به تغییرات محسوسی در زمان اجرا و Precision شده، اما سایر شاخص‌ها مانند Recall و ACC همچنان ثابت مانده‌اند که این نشان‌دهنده انعطاف‌پذیری و پایداری مدل پیشنهادی است. به‌طور کلی، جدول (۱۷) بیانگر این است که مدل پیشنهادی توانسته‌است در محیط‌های مقیاس‌پذیر، عملکردی مطلوب ارائه دهد و قابلیت استفاده در سامانه‌های آموزش الکترونیکی واقعی را دارد.

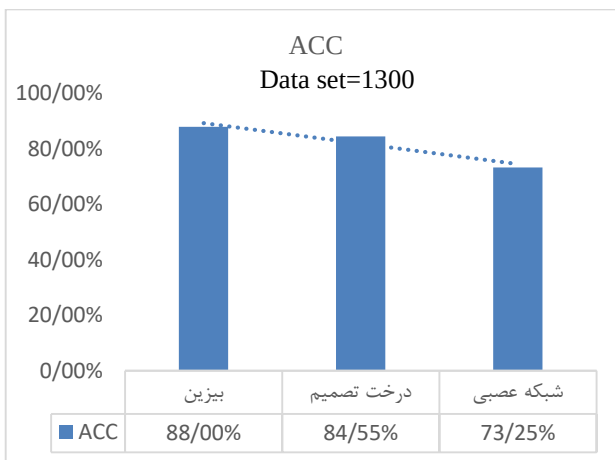
۱,۶ زمان اجرا

در این سناریو مدت زمان اجرای برنامه مورد ارزیابی قرار گرفته‌است. هر قدر زمان اجرا کمتر باشد الگوریتم کیفیت بالاتری دارد. در شکل (۳) زمان اجرای شبیه‌سازی آمده‌است.

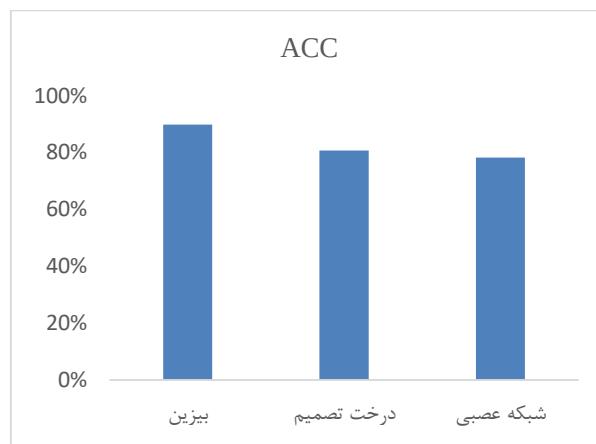


شکل ۳: مدت زمان اجرا

توانسته دقت، پایداری و عملکرد روش پیشنهادی را در مقایسه با درخت تصمیم و شبکه‌های عصبی به‌طور قابل توجهی بهبود بخشد.



شکل ۴: میزان دقت روش‌های مورد مقایسه برای ۱۳۰۰ کاربر



شکل ۵: میانگین دقت روش‌ها

۷. نتیجه‌گیری و کارهای آینده

- این پژوهش چارچوبی مبتنی بر ترکیب الگوریتم‌های K-Relief و Means و شبکه‌های بیزین ارائه کرد که با هدف بهبود عملکرد سیستم‌های آموزشی مبتنی بر وب طراحی شده بود. روش پیشنهادی تلاش کرد تا با شناسایی دقیق ویژگی‌های مهم، خوشه‌بندی کاربران مشابه و تحلیل روابط پیچیده میان داده‌ها، محدودیت‌های روش‌های سنتی را کاهش دهد و دقت پیش‌بینی و سرعت ارائه خدمات را افزایش دهد. نتایج آزمایش‌ها نشان داد که این چارچوب می‌تواند به‌طور مؤثری چالش‌هایی مانند آغاز سرد، مدیریت داده‌های ناقص و ارائه پیشنهادهای شخصی‌سازی شده را حل کند. در این پژوهش نکات برجسته‌ای وجود دارد که از جمله آن می‌توان به دقت بالا، روش پیشنهادی توانست با دقت پیش‌بینی متوسط ۹۲ درصد، به‌طور قابل توجهی از الگوریتم‌های رایج مانند درخت تصمیم (با دقت ۸۸ درصد) و شبکه عصبی (با دقت ۷۷ درصد) پیشی بگیرد. سرعت اجرا: زمان اجرای روش پیشنهادی با

بهره‌گیری از خوشه‌بندی K-Means و کاهش ابعاد داده‌ها، در تمام آزمایش‌ها در محدوده بهینه باقی‌ماند و مقیاس‌پذیری نسبی آن در برابر افزایش تعداد کاربران و خدمات ثابت شد. کاهش خطاها: کاهش نرخ خطاهای مثبت کاذب و منفی کاذب به ترتیب به میزان ۵٪ و ۴٪ در مقایسه با روش‌های مرجع، نشان‌دهنده کیفیت بالای پیشنهادهای ارائه‌شده به کاربران بود. غلبه بر آغاز سرد: خوشه‌بندی کاربران بر اساس ویژگی‌های استخراج‌شده، امکان ارائه پیشنهادهای منطقی حتی در صورت کمبود داده‌های تاریخی برای کاربران جدید را فراهم کرد. انعطاف‌پذیری در مدیریت داده‌های ناقص: شبکه بیزین توانست با تحلیل احتمالی، از تأثیر منفی داده‌های ناقص بر عملکرد سیستم بکاهد، اشاره کرد.

محدودیت‌ها و چالش‌هایی در پژوهش ما وجود که با برطرف کردن آن‌ها می‌توان تحقیقات جامع‌تری انجام داد از جمله این موارد:

- حساسیت به کیفیت داده‌های ورودی: روش انتخاب ویژگی و خوشه‌بندی به‌شدت به کیفیت داده‌های اولیه وابسته بود. داده‌های نویزی می‌توانند بر نتایج تأثیر منفی بگذارند.
- عدم پشتیبانی از تحلیل‌های بلادرنگ: روش ارائه‌شده قادر به تطبیق سریع با تغییرات آنی در رفتار کاربران نبود و نیازمند دوره‌های مشخص برای به‌روزرسانی مدل بود.
- کمبود آزمایش‌های عملی در محیط‌های واقعی: آزمایش‌ها عمدتاً بر داده‌های شبیه‌سازی شده انجام شدند و اثربخشی مدل در سناریوهای واقعی و متغیر هنوز به‌طور کامل ارزیابی نشده است. برای غلبه بر محدودیت‌های شناسایی شده و تقویت کارایی سیستم، پیشنهاد می‌شود محورهای زیر مورد بررسی و توسعه قرار گیرند:
- ادغام الگوریتم‌های فراابتکاری: استفاده از الگوریتم‌های فراابتکاری مانند ژنتیک، بهینه‌سازی ازدحام ذرات و کلونی مورچگان برای بهبود فرآیند انتخاب ویژگی و خوشه‌بندی کاربران. این روش‌ها می‌توانند با جستجوی راه‌حل‌های بهینه‌تر، تأثیر داده‌های ناقص را کاهش دهند و دقت پیش‌بینی را افزایش دهند.
- توسعه یادگیری عمیق: طراحی مدل‌هایی که از شبکه‌های عصبی کانولوشن (CNN) یا بازگشتی (RNN) در کنار شبکه‌های بیزین استفاده کنند، می‌تواند تحلیل الگوهای پیچیده‌تر را امکان‌پذیر سازد. این ترکیب به‌ویژه در تحلیل داده‌های غیرساختاریافته یا چندرسانه‌ای مفید خواهد بود.
- تحلیل بلادرنگ رفتار کاربران: توسعه سیستمی که بتواند تغییرات لحظه‌ای در رفتار کاربران را شناسایی و مدل را به‌صورت پویا به‌روزرسانی کند. این هدف می‌تواند با استفاده از روش‌های یادگیری آنلاین یا تطبیقی محقق شود.
- مدیریت داده‌های چندمنظوره: یکپارچه‌سازی داده‌های مختلف از منابع گوناگون مانند برنامه‌های تلفن همراه، وبسایت‌ها، و حسگرها برای ایجاد نمای کاملی از رفتار کاربران و بهبود شخصی‌سازی.

- [12] Z. Mousavi and M. Khan Babayi, "A data mining-based method for guiding students in educational systems," in 1st National Conference on Business Management, 2013.
- [13] R.B. Kaliwal, S.L. Deshpande. Bayesian Network Model Based Classifiers Are Used in an Intelligent E-learning System. In: Garg, D., Rodrigues, J.J.P.C., Gupta, S.K., Cheng, X., Sarao, P., Patel, G.S. (eds) Advanced Computing. IACC 2023. Communications in Computer and Information Science, vol 2053, 2024.
- [14] C. He, W. Liu and J. Ren, "Bayesian Network Structure Learning: A Review," 2022 6th Asian Conference on Artificial Intelligence Technology (ACAIT), Changzhou, China, pp. 1-7, 2022.
- [15] Li, G., & Wang, Y. Research on learner's emotion recognition for intelligent education system. In 2018 IEEE 3rd Advanced Information Technology, Electronic and Automation Control Conference (IAEAC), pp. 754-758, October 2018.
- [16] Qiu, D. Development and Implementation of Learning System of an Intelligent Learning System for Ideological and Political Education in Colleges under Mobile Platform. In 2018 International Conference on Virtual Reality and Intelligent Systems (ICVRIS), pp. 289-292, August 2018.
- [17] H.T. Kahraman, S. Sagioglu and I. Colak. Development of adaptive and intelligent web-based educational systems. In 2010 4th International Conference on Application of Information and Communication Technologies, pp. 1-5, 2010, October.
- [18] Hu, H., Li, J., Lei, X., Qin, P., & Chen, Q. Design of health statistics intelligent education system based on Internet+. In Journal of Physics: Conference Series (Vol. 1168, No. 6, p. 062003, February 2019.
- [19] Xu, Y. Physical Education Evaluation System Analysis and Intelligent Evaluation System Design. Caribbean Journal of Science, 52(4), pp. 1521-1524, 2019.
- [20] R. O'Brien, M. Hartnett, and P. Rawlins, (2019). The centralization of elearning resource development within the New Zealand vocational tertiary education sector. Australasian Journal of Educational Technology, 2019.
- [21] H. Bahador, S. Aliaei, and A. Bagheri, "Presenting a model for identifying factors affecting students' academic guidance using decision tree and neural network," in 2nd National Congress of New Technologies of Iran with the Aim of Achieving Sustainable Development, Tehran, Center for Achieving Sustainable Development Solutions, Mehr Arvand Higher Education Institute, 2015.
- [22] B. Maghsoudi, S. Soleimani, A. Amiri, and M. Afsharchi, "Improving the quality of education in e-learning systems using educational data mining," Journal of Educational Technology, vol. 6, 2012.
- [23] B. Minae, S. S. Mirfazel, and S. H. Hani, "Identifying factors affecting academic failure using association rules and cluster analysis," in 6th International Conference on Data Mining of Iran, 2013.
- [24] Z. S. Mousavi and M. Khanbabaei, "Presenting a method based on a hybrid data mining technique for students' academic guidance," in 1st National Conference on Business Management, Hamedan, Tolou Farzin Science and Industry Company, Bu-Ali Sina University, 2013.
- [25] F. Gharebaghi and A. Amiri, "A Novel Semi-Supervised Approach to Improve E-learning Performance," *Intelligent Multimedia Processing and Communication Systems*, vol. 3, no. 3, pp. 53-63, 2022.
- [26] N. Farghzadeh and N. Modiri, "Developing a Performance-Aware Cloud Infrastructure in E-Learning and Virtual Systems," *Intelligent Multimedia Processing and*
- توسعه مقیاس‌پذیری: بررسی و بهینه‌سازی زمان اجرا و کارایی در محیط‌های مقیاس بزرگ با استفاده از تکنیک‌های موازی‌سازی و رایانش ابری. این روش‌ها می‌توانند هزینه محاسباتی را کاهش داده و سیستم را برای حجم‌های بزرگ داده قابل‌اجرا کنند.
- مدیریت عدم قطعیت: استفاده از مدل‌های فازی در کنار شبکه‌های بیزین برای بهبود تحلیل در شرایطی که داده‌ها مبهم یا ناقص هستند. پیاده‌سازی این پیشنهادها می‌تواند چارچوب پیشنهادی را به سطح جدیدی از کارایی، پایداری و کاربردپذیری در سیستم‌های آموزشی مبتنی بر وب ارتقا دهد. این مسیر توسعه به افزایش رضایت کاربران، کاهش هزینه‌های عملیاتی و تقویت قابلیت‌های شخصی‌سازی منجر خواهد شد.

References

- [1] M.I. Burhan, A. Kumala Jaya, and L. Adira. "Bayesian Intelligent Tutoring System for Vocational High Schools." PENA TEKNIK: Jurnal Ilmiah Ilmu-Ilmu Teknik 9, no. 1, pp. 1-13, 2024.
- [2] M. Tao, "Application of Bayesian Model in the Construction of Personalized Learning Platform for Internet Education," 2024 International Conference on Distributed Computing and Optimization Techniques (ICDCOT), Bengaluru, India, pp. 1-5, 2024.
- [3] M. Chanthiran, A. B. Ibrahim, M. H. Abdul Rahman, and P. Mariappan, "Bayesian Network Approach in Educational Application Development: A Systematic Literature Review and Bibliometric Meta-Analysis", *International Journal of Artificial Intelligence*, vol. 9, no. 1, pp. 8-16, Jun. 2022.
- [4] R. B. Kaliwal and S. L. Deshpande, "Assessment Study For E-Learning Using Bayesian Network," 2021 International Conference on Artificial Intelligence and Smart Systems (ICAIS), Coimbatore, India, pp. 79-84, 2021.
- [5] I. Nekhaev, I. Zhuykov, S. Manukyants, Maslennikov, A. (2020). Applying Bayesian Network to Assess the Levels of Skills Mastering in Adaptive Dynamic OER-Systems. In: Silhavy, R., Silhavy, P., Prokopova, Z. (eds) *Software Engineering Perspectives in Intelligent Systems*. CoMeSySo 2020. *Advances in Intelligent Systems and Computing*, vol 1294, 2020.
- [6] G. Li and Y. Wang, "Learner's emotion recognition for intelligent education systems," in 3rd Adv. Inf. Technol. Electron. Automat. Control Conf., IEEE, pp. 754-758, 2018.
- [7] D. Qiu, "Development and implementation of an intelligent learning system for ideological and political education," in *Int. Conf. on Virtual Reality and Intelligent Systems*, IEEE, pp. 289-292, 2018.
- [8] Y. Zhang, "Exploration on teaching reform of architectural design course in higher vocational colleges in the era of intelligent education," *Advances in Higher Education*, vol. 3, no. 3, pp. 33-35, 2019.
- [9] M. Ally, "Foundations of educational theory for online learning," in *Theory and Practice of Online Learning*, T. Anderson, Ed., Edmonton: AU Press, pp. 15-44, 2008.
- [10] F. Canales et al., "Adaptive and intelligent web-based education system: Towards an integral architecture and framework," *Expert Systems with Applications*, vol. 33, no. 4, pp. 1076-1089, 2007.
- [11] H. T. Kahraman, S. Sagioglu, and I. Colak, "Development of adaptive and intelligent web-based educational systems," in 4th Int. Conf. on Application of Information and Communication Technologies, IEEE, pp. 1-5, 2010.

Communication Systems, vol. 2, no. 1, pp. 53-64, Mar. 2021.

- [27] M. Bazargani, S. H. Alizadeh, and B. Masoumi, "Group deep neural network approach in semantic recommendation system for movie recommendation in online networks," *Electron. Commer. Res.*, 2024.