



ORIGINAL RESEARCH PAPER

Improving the accuracy of predicting students' academic success and failure based on the efficiency of deep feedforward neural networks

H. Ziafat^{*1}, Gh. Azariarani²

¹ Department of Computer Engineering, Nat. C., Islamic Azad University, Natanz, Iran

² Department of Computer Engineering, Faculty of Electrical and Computer Engineering, Technical and Vocational University, (TVU), Tehran, Iran

ABSTRACT

Received: 12 September 2025
Reviewed: 25 October 2025
Revised: 14 December 2025
Accepted: 30 January 2026

KEYWORDS:

Academic Performance Prediction
Educational Data Mining
Extreme Learning Machine
Relief Feature Selection

* Corresponding author

✉ Dr.ziafat@iau.ac.ir

☎ (+98913) 1639308

Background and Objective: Educational data mining, as a modern and interdisciplinary field, plays a crucial role in analyzing learners' behavior, identifying factors influencing academic success or failure, and supporting educational decision-making. With the growth of e-learning systems and the increasing volume of educational data, the need for intelligent methods to extract hidden knowledge from these data has become more pronounced. Although traditional machine learning methods have been widely used in numerous studies, they face limitations in terms of accuracy and generalization capability when dealing with complex, nonlinear, and high-dimensional educational data. In this context, deep learning-based models and artificial neural networks have attracted significant attention due to their strong ability to model complex relationships. The main objective of this study is to develop and evaluate an intelligent hybrid model based on feedforward neural networks, including Artificial Neural Networks (ANN), Extreme Learning Machines (ELM), and Multilayer Perceptron (MLP), combined with advanced feature selection methods, to identify key variables affecting students' academic performance and significantly improve the accuracy of academic status prediction.

Methods: In this study, a comprehensive set of educational data was utilized, including demographic, socio-economic, enrollment-related information, educational records, and students' academic grades. The data were collected from university educational systems as well as reputable public datasets. After performing necessary preprocessing steps such as data cleaning, handling missing values, and feature normalization, the feature selection process was carried out using three methods: MRMR, Chi-square, and ReliefF. Subsequently, the three neural networks—ANN, ELM, and MLP—were trained using 90% of the data, while the remaining 10% was used for testing and model validation. Finally, the outputs of the networks were integrated into a hybrid model using a majority voting strategy. The performance of the proposed model was evaluated using standard metrics, including accuracy, precision, recall, and the F-measure.

Findings: The experimental results demonstrated that the ReliefF method outperformed the other feature selection techniques in identifying influential features. Using the top 20 features selected by this method, the proposed hybrid model achieved an accuracy of 81.44% and an F-measure of 72.09%. In the evaluation of individual neural networks, the ELM model exhibited the best performance with an accuracy of 82.8%, which was on average 2–4% higher than that of ANN and MLP. Moreover, comparison with traditional machine learning approaches revealed that the hybrid neural network model improved prediction accuracy by more than 7% and classification accuracy by over 4%, indicating the significant superiority of the proposed approach.

Conclusion: The findings of this study indicate that combining feedforward neural networks with appropriate feature selection methods provides an efficient and reliable approach for predicting students' academic status. The proposed model can serve as an intelligent decision-support tool in educational systems to enable early identification of students at risk of academic failure. Future research may focus on incorporating attention mechanisms, more advanced feature selection techniques, deeper learning models, and larger and more diverse datasets to further enhance the accuracy and generalizability of the model.



COPYRIGHTS

© 2026 The Author(s). This is an open-access article distributed under the terms and conditions of the Creative Attribution-NonCommercial 4.0 International (CC BY-NC 4.0)
(<https://creativecommons.org/licenses/by-nc/4.0/>)



NUMBER OF REFERENCES

38



NUMBER OF FIGURES

17



NUMBER OF TABLES

1

مقاله پژوهشی

بهبود دقت پیش‌بینی افت و موفقیت تحصیلی دانشجویان بر پایه بهره‌وری از شبکه‌های عصبی پیش‌خور عمیق

حسن ضیافت^{۱*}، قاسم آذری آرانی^۲^۱گروه مهندسی کامپیوتر، واحد نطنز، دانشگاه آزاد اسلامی، نطنز، ایران^۲گروه آموزشی مهندسی کامپیوتر، دانشکده مهندسی برق و کامپیوتر، دانشگاه ملی مهارت، تهران، ایران

چکیده

پیشینه و اهداف: داده‌کاوی آموزشی به‌عنوان یکی از حوزه‌های نوین و میان‌رشته‌ای، نقش مهمی در تحلیل رفتار فراگیران، شناسایی عوامل مؤثر بر موفقیت یا افت تحصیلی و پشتیبانی از تصمیم‌گیری‌های آموزشی ایفا می‌کند. با رشد سامانه‌های آموزش الکترونیکی و افزایش حجم داده‌های آموزشی، نیاز به روش‌های هوشمند برای استخراج دانش نهفته در این داده‌ها بیش‌ازپیش احساس می‌شود. روش‌های سنتی یادگیری ماشین، هرچند در بسیاری از پژوهش‌ها به کار گرفته شده‌اند، اما در مواجهه با داده‌های پیچیده، غیرخطی و چندبعدی آموزشی با محدودیت‌هایی از نظر دقت و قدرت تعمیم‌پذیری روبه‌رو هستند. در این میان، مدل‌های مبتنی بر یادگیری عمیق و شبکه‌های عصبی مصنوعی به دلیل توانایی بالا در مدل‌سازی روابط پیچیده، توجه پژوهشگران را به خود جلب کرده‌اند. هدف اصلی این پژوهش، توسعه و ارزیابی یک مدل ترکیبی هوشمند مبتنی بر شبکه‌های عصبی پیش‌خور شامل شبکه عصبی مصنوعی، شبکه یادگیری با سرعت بالا (ELM) و شبکه پرسپترون چندلایه به همراه روش‌های پیشرفته انتخاب ویژگی است تا ضمن شناسایی متغیرهای کلیدی مؤثر بر عملکرد تحصیلی دانشجویان، دقت پیش‌بینی وضعیت تحصیلی آنان به طور معناداری افزایش یابد.

روش‌ها: در این مطالعه، مجموعه‌ای جامع از داده‌های آموزشی شامل اطلاعات دموگرافیک، اجتماعی-اقتصادی، وضعیت ثبت‌نام، سوابق آموزشی و نمرات تحصیلی دانشجویان مورد استفاده قرار گرفت. داده‌ها از سامانه‌های آموزشی دانشگاهی و همچنین مجموعه داده‌های عمومی معتبر گردآوری شدند. پس از انجام پیش‌پردازش‌های لازم از جمله پاک‌سازی داده‌ها، مدیریت مقادیر گمشده و نرمال‌سازی ویژگی‌ها، فرایند انتخاب ویژگی با استفاده از سه روش MRMR، Chi-square و ReliefF انجام شد. سپس سه شبکه عصبی ANN، ELM و MLP با استفاده از ۹۰ درصد داده‌ها آموزش داده شدند و ۱۰ درصد باقی‌مانده برای آزمون و اعتبارسنجی مدل‌ها مورد استفاده قرار گرفت. در نهایت، خروجی شبکه‌ها با بهره‌گیری از روش رأی‌گیری اکثریت به‌صورت یک مدل ترکیبی ادغام شد. عملکرد مدل پیشنهادی با استفاده از معیارهای استاندارد دقت، صحت، فراخوانی و معیار F ارزیابی گردید.

یافته‌ها: نتایج حاصل از آزمایش‌ها نشان داد که روش ReliefF در مقایسه با سایر روش‌های انتخاب ویژگی، عملکرد بهتری در شناسایی ویژگی‌های مؤثر دارد. با استفاده از ۲۰ ویژگی برتر استخراج‌شده توسط این روش، مدل ترکیبی پیشنهادی به دقت ۸۱.۴۴ درصد و مقدار F برابر با ۷۲/۰۹ درصد دست‌یافت. در بررسی عملکرد شبکه‌های عصبی به‌صورت مستقل، شبکه ELM بهترین نتیجه را با دقت ۸۲/۸ درصد ارائه کرد که به طور متوسط ۲ تا ۴ درصد بالاتر از ANN و MLP بود. علاوه بر این، مقایسه مدل پیشنهادی با روش‌های سنتی یادگیری ماشین نشان داد که مدل ترکیبی شبکه‌های عصبی توانسته است دقت پیش‌بینی را بیش از ۷ درصد و دقت طبقه‌بندی را بیش از ۴ درصد بهبود بخشد که بیانگر برتری معنادار رویکرد پیشنهادی است.

تاریخ دریافت: ۲۱ شهریور ۱۴۰۴
تاریخ دوری: ۰۳ آبان ۱۴۰۴
تاریخ اصلاح: ۲۳ آذر ۱۴۰۴
تاریخ پذیرش: ۱۰ بهمن ۱۴۰۴

واژگان کلیدی:

پیش‌بینی وضعیت تحصیلی
داده‌کاوی آموزشی
ماشین یادگیری افراطی
الگوریتم Relief

* نویسنده مسئول

Dr.ziafat@iaua.ac.ir

① ۰۹۱۳-۱۶۳۹۳۰۸

نتیجه گیری: یافته‌های این پژوهش نشان می‌دهد که ترکیب شبکه‌های عصبی پیش‌خور با روش‌های مناسب انتخاب ویژگی، رویکردی کارآمد و قابل‌اتکا برای پیش‌بینی وضعیت تحصیلی دانشجویان است. مدل پیشنهادی می‌تواند به‌عنوان یک ابزار تصمیم‌یار هوشمند در نظام‌های آموزشی برای شناسایی زودهنگام دانشجویان در معرض افت تحصیلی مورد استفاده قرار گیرد. در پژوهش‌های آینده، به‌کارگیری سازوکارهای توجه، روش‌های پیشرفته‌تر انتخاب ویژگی، مدل‌های یادگیری عمیق‌تر و استفاده از مجموعه داده‌های بزرگ‌تر و متنوع‌تر می‌تواند به بهبود بیشتر دقت و تعمیم‌پذیری مدل منجر شود.

مقدمه

جدیدی را بر اساس الگوریتم‌های یادگیری ماشینی برای پیش‌بینی نمرات آزمون نهایی دانشجویان مقطع کارشناسی، با در نظر گرفتن نمرات آزمون‌های میان‌ترم به‌عنوان داده‌های منبع، پیشنهاد می‌کند. عملکرد جنگل‌های تصادفی، ماشین‌های بردار پشتیبان، رگرسیون لجستیک، الگوریتم‌های ساده بیز و k-نزدیک‌ترین همسایه که از جمله الگوریتم‌های یادگیری ماشینی هستند، برای پیش‌بینی نمرات آزمون نهایی مقایسه شد. یکی از دلایل عدم موفقیت چنین مدل‌هایی اکتفا کردن به روش‌های یادگیری ماشینی سنتی است در صورتی که اغلب پژوهش‌های موجود نشان داده‌اند که شبکه‌های عصبی و یادگیری عمیق عموماً نتایج امیدوارکننده‌تری در مقایسه با روش‌های یادگیری ماشینی سنتی ارائه می‌دهند.

در این راستا، در این مقاله یک مدل مبتنی بر شبکه‌های عصبی پیش‌خور عمیق شامل شبکه عصبی چندلایه پرسپترون و شبکه عصبی ماشین یادگیری افراطی و شبکه عصبی پیش‌خور ساده ارائه شده است که هدف آن پیش‌بینی وضعیت تحصیلی دانش‌آموزان است. در این مدل یک مرحله انتخاب ویژگی درج می‌شود تا از میان ویژگی‌های مختلف آن دسته از ویژگی‌هایی که بر عملکرد تحصیلی دانش‌آموزان مؤثر است شناسایی شود. برای بهبود دقت پیش‌بینی وضعیت دانش‌آموزان از ترکیب سه شبکه عصبی فوق و رأی‌گیری اکثریت استفاده می‌شود تا طبقه‌بندی نهایی انجام شود. براین اساس، اهداف اصلی این پژوهش عبارت‌اند از:

- طراحی و ارائه یک مدل ترکیبی مبتنی بر شبکه‌های عصبی پیش‌خور شامل شبکه عصبی پیش‌خور ساده (ANN)، ماشین یادگیری افراطی (ELM) و شبکه عصبی عمیق چندلایه پرسپترون (MLP) به‌منظور پیش‌بینی وضعیت تحصیلی دانشجویان؛

- بررسی و مقایسه عملکرد سه روش انتخاب ویژگی مبتنی بر فیلتر شامل Chi-square، ReliefF و MRMR در بهبود دقت پیش‌بینی؛

- تحلیل تأثیر تعداد ویژگی‌های منتخب بر عملکرد مدل پیشنهادی؛

- ارزیابی و مقایسه عملکرد شبکه‌های عصبی مورد استفاده و مدل ترکیبی مبتنی بر رأی‌گیری اکثریت از منظر معیارهای دقت، صحت، فراخوان و معیار F؛

- مقایسه نتایج مدل پیشنهادی با روش‌های یادگیری ماشینی سنتی گزارش شده در مطالعات پیشین.

در ادامه این مقاله در بخش ۲ شبکه‌های عصبی بررسی می‌شوند، و مطالعات این حوزه معرفی می‌شوند. در بخش ۳ روش پیشنهادی

داده‌های جمع‌آوری شده در مورد فرایندهای آموزشی فرصت‌های جدیدی را برای بهبود تجربه یادگیری و بهینه‌سازی تعامل کاربران با پلتفرم‌های فناوری ارائه می‌دهد [۱]. پردازش داده‌های آموزشی باعث بهبود در بسیاری از زمینه‌ها مانند پیش‌بینی رفتار دانش‌آموزان، یادگیری تحلیلی و رویکردهای جدید در سیاست‌گذاری آموزشی می‌شود [۲، ۳]. این مجموعه جامع از داده‌ها نه تنها به مقامات آموزشی اجازه می‌دهد تا سیاست‌های مناسبی مبتنی بر داده اتخاذ کنند، بلکه اساس نرم‌افزاری را تشکیل می‌دهد که باهوش مصنوعی در فرایند یادگیری توسعه می‌یابد. داده‌کاوی آموزشی مربیان را قادر می‌سازد تا موقعیت‌هایی مانند ترک تحصیل یا علاقه کمتر به درس را پیش‌بینی کنند، عوامل درونی مؤثر بر عملکرد آنها را تجزیه و تحلیل کنند و روش‌های آماری را برای پیش‌بینی عملکرد تحصیلی دانش‌آموزان ارائه دهند. انواع روش‌های داده‌کاوی برای پیش‌بینی عملکرد دانش‌آموزان، شناسایی دانش‌آموزان ضعیف و ترک تحصیل استفاده می‌شود [۳]. پیش‌بینی زودهنگام پدیده جدیدی است که شامل روش‌های ارزیابی برای حمایت از دانش‌آموزان با پیشنهاد راهبردها و سیاست‌های اصلاحی مناسب در این زمینه است [۴]. اکنون دانشگاه‌ها باید ظرفیت استفاده از این داده‌ها را برای پیش‌بینی موفقیت تحصیلی و تضمین پیشرفت دانش‌آموزان بهبود بخشند [۵]. در نتیجه، داده‌کاوی آموزشی با کشف الگوهای پنهان در داده‌های آموزشی، اطلاعات جدیدی را در اختیار مربیان قرار می‌دهد. با استفاده از این مدل می‌توان برخی از جنبه‌های نظام آموزشی را ارزیابی و بهبود بخشید تا از کیفیت آموزش اطمینان حاصل شود.

به دلیل مزایای بسیاری که مؤسسات آموزشی می‌توانند به آن دست یابند، داده‌کاوی آموزشی که شامل آموزش و انفورماتیک است به یک حوزه تحقیقاتی ضروری تبدیل شده است. در امتداد این خطوط، روش‌های مختلف داده‌کاوی برای بهبود نتایج یادگیری با کاوش در داده‌های مقیاس بزرگ که از محیط‌های آموزشی به دست می‌آیند، استفاده شده‌اند. یکی از مشکلات اصلی پیش‌بینی دستاوردهای آینده دانش‌آموزان است؛ با تخمین دستاوردهای آینده دانش‌آموزان می‌توان فعالانه به دانش‌آموزان در رسیدن به عملکرد بهتر و جلوگیری از ترک تحصیل کمک کرد. تلاش‌های زیادی برای حل مشکل پیش‌بینی عملکرد دانش‌آموزان در زمینه داده‌کاوی آموزشی صورت گرفته است [۶]. یکی از آخرین تلاش‌های موجود در این زمینه مربوط به مطالعه مصطفی یا گچی و همکاران در سال ۲۰۲۲ [۷] است. این مطالعه مدل

پیش‌پردازش پیچیده و هزینه محاسباتی بالا از محدودیت‌های این روش گزارش شده است.

در [۱۱] به بررسی روش‌های مختلف داده‌کاوی در راستای پیش‌بینی سطح عملکرد دانشجویان پرداخته شده است. بدین منظور از مجموعه داده‌های کال بورد ۳۶۰ و به‌منظور تحلیل‌ها و بررسی‌ها از ۵ روش داده‌کاوی متفاوت بهره برده شده است. مجموعه داده‌های استفاده شده شامل صد و شصت و سه نمونه و شانزده ویژگی است. در پایان نتایج حاکی از آن است که روش شبکه عصبی چندلایه پرسپترون با میزان صحت ۷۶٪ برترین واکنش را در قیاس با سایر روش‌ها داشته است.

در [۱۲] به بحث داده‌کاوی آموزشی برای تعیین موفقیت دانشجویان در مؤسسات آموزش عالی آفریقای جنوبی پرداخته شده است. در این پژوهش، روش‌های طبقه‌بندی مختلفی برای یک مجموعه داده مصنوعی تولید شده توسط یک شبکه بیزی اعمال شده تا نشان دهد این طبقه‌بندی‌کننده‌ها می‌توانند برای پیش‌بینی عملکرد دانش‌آموز از قبل باهدف ارتقای موفقیت دانش‌آموز و جلوگیری از پیامدهای منفی ناشی از تلاش دانشجویان برای تکمیل تحصیل یا ترک تحصیل استفاده شوند. در [۱۳] عوامل متعددی که باعث ترک تحصیل دانشجویان شده بررسی شده و برای این کار بر پیش‌بینی دانش‌آموزانی که در معرض خطر ترک تحصیل هستند با استفاده از پنج الگوریتم طبقه‌بندی K نزدیک‌ترین همسایه، بیزین ساده، درخت تصمیم، جنگل تصادفی و ماشین بردار پشتیبان تمرکز شده است. نتایج این تحقیق می‌تواند به بهبود شیوه‌های آموزشی به‌منظور بهبود عملکرد دانش‌آموزانی که در معرض خطر ترک تحصیل هستند، کمک شایان توجهی نماید.

در [۱۴] به نگرانی‌ها در ارتباط با ترک تحصیل دانشجویان مؤسسات آموزش عالی پرداختند؛ زیرا این امر تأثیر زیادی نه تنها بر هزینه‌های دانشجویان، بلکه بر هدررفتن بودجه عمومی دارد. در این کار، درخت تصمیم، رگرسیون لجستیک، جنگل تصادفی، K نزدیک‌ترین همسایه و الگوریتم شبکه عصبی برای ارائه بهترین مدل مقایسه شده که نشان می‌دهد مدل رگرسیون لجستیک بهترین یادگیرنده برای پیش‌بینی ترک تحصیل دانشجویان با علل موضوعی بالقوه است.

در [۱۵] به بررسی حوزه پیش‌بینی عملکرد دانش‌آموزان با بهره‌وری از روش‌های مبتنی بر هوش مصنوعی و داده‌کاوی آموزشی پرداختند. این کار مروری نظام‌مند از مطالعه پیش‌بینی عملکرد دانش‌آموزان از دیدگاه یادگیری ماشین و داده‌کاوی ارائه می‌نماید. برای درک شهودی آزمایش‌هایی روی یک مجموعه داده آموزشی متشکل از ۱۳۲۵ دانشجوی و ۸۳۲ دوره از سامانه اطلاعاتی انجام شده است. نتایج این تحقیق پیشرفت‌ها و چالش‌های مرتبط با پیش‌بینی عملکرد دانش‌آموزان را به‌خوبی فراهم نموده و در اختیار سایر محققین قرار داده است.

در [۱۶] به بحث پیش‌بینی عملکرد دانش‌آموزان در امتحانات به واسطه درخت تصمیم و K نزدیک‌ترین همسایه پرداخته شده است. نتایج این مطالعه نشان می‌دهد که درخت تصمیم برای پیش‌بینی وضعیت

معرفی می‌شود و در بخش ۴ نتایج ارزیابی روش ارائه می‌شوند و در بخش پایانی نتیجه‌گیری ارائه خواهد شد.

پژوهش‌های مرتبط

باتوجه به مباحثی که در ارتباط با اهمیت بحث پیش‌بینی وضعیت تحصیلی دانش‌آموزان و داده‌کاوی آموزشی ارائه و عنوان گردید، تاکنون تحقیقات گسترده‌ای در این ارتباط پیشنهاد شده‌اند. در این بخش تعدادی از مهم‌ترین این تحقیقات را در قالب پیشینه پژوهش معرفی نموده و از دیدگاه‌های مختلف آن‌ها را ارزیابی و تحلیل خواهیم نمود.

در [۸] یک چارچوب نوین برای پیش‌بینی زودهنگام عملکرد تحصیلی دانشجویان با بهره‌گیری از مدل‌های زبانی بزرگ ارائه شده است. روش پیشنهادی با ترکیب داده‌های ساخت‌یافته آموزشی (نمرات، حضور، فعالیت‌های یادگیری) و داده‌های متنی غیرساخت‌یافته، از قابلیت درک معنایی LLMها برای استخراج ویژگی‌های غنی استفاده می‌کند. مدل LLM-EPSP قادر است الگوهای پنهان رفتاری و آموزشی دانشجویان را در مراحل ابتدایی دوره شناسایی کند. نتایج تجربی نشان می‌دهد که این روش در مقایسه با مدل‌های یادگیری ماشین و شبکه‌های عصبی کلاسیک، بهبود معناداری در دقت و F-measure پیش‌بینی زودهنگام عملکرد تحصیلی ارائه می‌دهد. یافته‌ها بیانگر پتانسیل بالای LLMها در داده‌کاوی آموزشی و سامانه‌های هشدار زودهنگام افت تحصیلی هستند. در [۹] مجموعه‌ای از مدل‌های یادگیری ماشین و یادگیری عمیق شامل CatBoost، LightGBM، شبکه عصبی مصنوعی، شبکه عصبی عمیق، شبکه عصبی عمیق وزن‌دار (WDNN) برای پیش‌بینی پیامدهای تحصیلی دانشجویان مورد مقایسه قرار گرفته‌اند. نوآوری اصلی این مطالعه، استفاده از منظم‌سازی ElasticNet به‌منظور پایش بیش‌برازش و بهبود پایداری مدل‌ها در مواجهه با داده‌های آموزشی پرنویز و هم‌بسته است. نتایج تجربی نشان می‌دهد که مدل‌های مبتنی بر گرادیان بوستینگ، به‌ویژه CatBoost و LGBM، در کنار WDNN منظم‌سازی شده، عملکرد بهتری نسبت به ANN و DNN کلاسیک ارائه می‌دهند. همچنین به‌کارگیری ElasticNet منجر به افزایش دقت پیش‌بینی و بهبود تعمیم‌پذیری مدل‌ها شده است. این مطالعه بر اهمیت ترکیب معماری‌های پیشرفته با روش‌های منظم‌سازی در تحلیل داده‌های آموزشی تأکید می‌کند.

در [۱۰] چارچوبی نوآورانه برای پیش‌بینی موفقیت تحصیلی دانشجویان با بهره‌گیری از یادگیری تجمعی و تبدیل داده‌های رفتاری به نمایش‌های تصویری ارائه می‌شود. در روش پیشنهادی، داده‌های رفتاری دانشجویان به تصاویر دوبعدی نگاشت شده و سپس با استفاده از مدل‌های یادگیری عمیق و الگوریتم‌های تجمعی مورد تحلیل قرار می‌گیرند. نتایج تجربی نشان می‌دهد که مدل‌های Ensemble مبتنی بر داده‌های تصویری، دقت پیش‌بینی بالاتری نسبت به روش‌های کلاسیک مبتنی بر ویژگی‌های عددی ارائه می‌کنند. این رویکرد توانایی مدل در استخراج الگوهای پیچیده رفتاری را به طور معناداری افزایش می‌دهد. بااین‌حال، نیاز به

گرفته است. برای این منظور از مجموعه داده از مخزن Kaggle استفاده گردیده و تجزیه و تحلیل مجموعه داده ها با استفاده از الگوریتم های طبقه بندی مختلف مانند درخت تصمیم، جنگل تصادفی، طبقه بندی کننده SVM، طبقه بندی کننده SGD (Stochastic Gradient Descent)، طبقه بندی کننده AdaBoost و طبقه بندی کننده Logistic (LR) (Regression) انجام شده است. نتایج نشان می دهد که جنگل تصادفی بالاترین امتیازی (۹۸٪) را کسب نموده است. در مجموع نتایج نشان می دهد که رسانه های اجتماعی به میزان زیادی بر عملکرد دانش آموزان تأثیر می گذارد.

در [۲۱] از روش های مختلف یادگیری ماشینی برای پیش بینی عملکرد تحصیلی دانش آموزان با استفاده از داده های واقعی جمع آوری شده (شامل تاریخچه تحصیلی و عادات شخصی دانش آموزان) استفاده نمودند. علاوه بر این، مقایسه ای از روش های یادگیری ماشینی بر روی معیارهای ارزیابی مختلف ارائه شده است. نتایج این تحقیق این به دانش آموزان کمک می کند تا عملکرد تحصیلی خود را پیگیری کرده و براین اساس، الگوی مطالعه خود را مدیریت کنند تا به آنها در آینده کمک کند.

در [۲۲] یک مدل ترکیبی دوبعدی شبکه عصبی کانولوشن به منظور پیش بینی عملکرد تحصیلی دانشجویان پیشنهاد شده است. به منظور ارزیابی از تبدیل داده ها یک بعدی به دوبعدی استفاده شده و نتایج ارزیابی بر این مهم دلالت دارد که مدل پیشنهادی در این کار در قیاس با سایر مدل ها (اعم از k نزدیک ترین همسایه، بیز ساده، درخت های تصمیم گیری و رگرسیون لجستیک) در قبال سناریوهای مختلف با دقت بالاتر و بهتر عمل نموده است.

در [۲۳] یک الگوریتم یادگیری ترکیبی جدید تحت عنوان HELA (novel hybrid ensemble learning algorithm) برای پیش بینی عملکرد تحصیلی دانش آموزان پیشنهاد شده است. عملکرد دانش آموزان در کلاس های ریاضی و پرتغالی توسط الگوریتم پیشنهادی پیش بینی شده و بر حسب نتایج تجربی، مقادیر دقت ۹۶/۶٪ و ۹۱/۲٪ به ترتیب برای درس ریاضیات و درس پرتغالی به دست آمده است.

در [۲۴] به مسئله عدم تعادل کلاس در مجموعه داده های آموزشی و تأثیرات آن بر افت شدید پیش بینی ها پرداخته شده است. در این کار چندین روش نمونه برداری برای رسیدگی به نسبت های مختلف مشکل عدم تعادل طبقاتی با استفاده از مجموعه داده های مطالعات طولی دبیرستانی مقایسه شده است. برای مقایسه، از نمونه برداری بیش از حد تصادفی، کم نمونه گیری تصادفی و ترکیبی از روش نمونه برداری اقلیت مصنوعی برای اسمی و پیوسته به عنوان یک روش نمونه گیری مجدد ترکیبی استفاده شده است. همچنین از جنگل تصادفی به عنوان الگوریتم طبقه بندی برای ارزیابی نتایج هر روش نمونه گیری استفاده گردیده است. نتایج نشان می دهد که نمونه برداری بیش از حد تصادفی برای داده های نسبتاً نامتعادل و نمونه گیری مجدد ترکیبی برای داده های بسیار نامتعادل بهترین کارایی را داشته اند.

قبولی/عدم شکست دانش آموزان در یک دوره تحصیلی موفق ترین نتایج را ارائه داده و نرخ موفقیت ۹۱ درصد را فراهم ساخته است. این مدل به مدرسین کمک می کند تا اقدامات لازم را برای کمک به دانشجویان به طور مؤثر انجام دهند.

در [۱۷] مدل جدیدی بر اساس الگوریتم های یادگیری ماشینی برای پیش بینی نمرات امتحان نهایی دانشجویان مقطع کارشناسی، با در نظر گرفتن نمرات امتحانات میان ترم، پیشنهاد شده و عملکرد الگوریتم های جنگل های تصادفی، ماشین بردار پشتیبان، رگرسیون لجستیک، الگوریتم های ساده بیز و k نزدیک ترین همسایه برای پیش بینی نمرات امتحان نهایی بررسی و مقایسه شده است. نتایج نشان می دهد که مدل پیشنهادی در این کار به دقت طبقه بندی ۷۵-۷۰ درصدی دست یافته است. در این کار، پیش بینی ها تنها با استفاده از سه نوع متغیر شامل نمرات امتحانات میان ترم، داده های گروه و داده های دانشکده انجام شده و از این نظر دقت در پیش بینی ها پایین است. اما با این حال چنین مطالعاتی می توانند به تصمیم گیری آموزشی و پیش بینی اولیه دانش آموزان در معرض خطر بالا کمک شایان توجهی نمایند.

در [۱۷] به مسئله توزیع نامتعادل طبقات در داده های جمع آوری شده به عنوان یکی از چالش های مهم ایجاد نظام های طبقه بندی و پیش بینی دقیق پرداخته شده است. این مطالعه به بررسی تکنیک های سطح داده و تکنیک های سطح الگوریتم می پردازد. شش طبقه بندی کننده از هر تکنیک برای بررسی اثربخشی آن ها برای رسیدگی به مشکل داده های نامتعادل در حین پیش بینی نمره فارغ التحصیلی دانش آموزان بر اساس عملکرد آنها در مرحله اول استفاده شده است. نتایج نشان می دهد طبقه بندی کننده ها با مجموعه داده های نامتعادل عملکرد خوبی نداشته و با استفاده از این تکنیک ها می توان عملکردها را بهبود بخشید.

در [۱۸] یک مدل جدید پیش بینی پیشرفت دانشجویان بر اساس شبکه عصبی Spiking تکاملی پیشنهاد شده است. الگوریتم غشای تکاملی برای یادگیری فرایارامترهای مدل معرفی گردیده تا دقت مدل را در پیش بینی پیشرفت دانش آموزان بهبود بخشد. نتایج تجربی نشان می دهد که مدل مبتنی بر شبکه عصبی اسپایکینگ می تواند به طور مؤثری دقت پیش بینی پیشرفت دانش آموزان را بهبود بخشد.

در [۱۹] روش های مورد استفاده برای پیش بینی عملکرد تحصیلی دانش آموزان از سال ۲۰۱۵ تا ۲۰۲۱ مورد بحث و بررسی قرار گرفتند. این مطالعه ۵۸ مقاله از ۲۱۹ مقاله پژوهشی را از پایگاه های داده لنز و اسکوپوس بررسی می کند. دستاوردهای مطالعات نشان می دهد که سوابق تحصیلی و جمعیت شناسی دانش آموزان عوامل اولیه ای هستند که بر عملکرد دانش آموزان تأثیر می گذارند. علاوه بر این، نتایج حاکی از آن است که از میان الگوریتم های داده کاوی، طبقه بندی کننده درخت تصمیم پرکاربردترین الگوریتم داده کاوی است.

در [۲۰] ارزیابی روش های داده کاوی برای تجزیه و تحلیل تأثیر رسانه های اجتماعی بر عملکرد تحصیلی دانش آموزان دبیرستانی مورد توجه قرار

ریاضیات استفاده نموده است. تکنیک‌های درخت تصمیم و بیز ساده به ترتیب بهترین عملکرد پیش‌بینی‌کننده را برای دروس انگلیسی و ریاضیات ارائه داده‌اند. این مطالعه عملکرد تحصیلی گذشته دانش‌آموزان را به‌عنوان مهم‌ترین پیش‌بینی‌کننده معرفی می‌نماید. جدول ۱ این پژوهش‌ها را با یکدیگر مقایسه می‌کند.

در [۲۵] پیش‌بینی عملکرد دانش‌آموزان در زبان انگلیسی و ریاضیات با استفاده از روش‌های داده‌کاوی موردتوجه قرار گرفته است. این مطالعه از داده‌های بایگانی شده دانش‌آموزانی استفاده نمود که در سال ۲۰۱۹ شانزده‌ساله بوده و در گواهینامه آزمون مالزی در سال ۲۰۲۱ شرکت کردند. این مطالعه از قواعد درخت تصمیم برای تعیین ویژگی‌های دانش‌آموزان با عملکرد پایین، متوسط و بالا در دروس انگلیسی و

جدول ۱: مقایسه پژوهش‌های پیشین
Table 1: Comparison of Previous Research

Weaknesses نقاط ضعف	Strengths نقاط قوت	Proposed Method روش پیشنهادی	Author & Year نام نویسنده و سال
Limited in modeling complex non-linear relationships محدودیت در مدل‌سازی روابط غیرخطی پیچیده	Use of standard ML methods, simple and interpretable استفاده از روش‌های استاندارد یادگیری ماشین، ساده و قابل تفسیر	Machine Learning Algorithms for Student Performance Prediction الگوریتم‌های یادگیری ماشین برای پیش‌بینی عملکرد دانشجویان	M. Yağcı, 2022
High computational resource requirements, model complexity نیاز به منابع محاسباتی بالا، پیچیدگی مدل	Ability to process textual data and early prediction, use of LLM قدرت پردازش داده‌های متنی و پیش‌بینی زودهنگام، استفاده از LLM	LLM-EPSP: Using Large Language Models for Early Prediction of Student Performance LLM-EPSP : استفاده از مدل‌های زبانی بزرگ برای پیش‌بینی زودهنگام عملکرد دانشجویان	H. Zhou et al., 2026
High complexity, requires extensive hyperparameter tuning پیچیدگی زیاد، نیاز به تنظیمات زیاد هایپرپارامتر	Comprehensive algorithm comparison, use of regularization to reduce overfitting مقایسه جامع الگوریتم‌ها، استفاده از منظم‌سازی برای کاهش overfitting	Comparison of ML and DL with ElasticNet: LGBM, CatBoost, ANN, DNN, and WDNN مقایسه ML و DL با ElasticNet: LGBM, CatBoost, ANN, DNN و WDNN	V. Eriyandi & A. N. Ahmad, 2026
Requires image data and heavy processing نیاز به داده‌های تصویری و پردازش سنگین	Combination of image and behavioral data, improved accuracy ترکیب داده‌های تصویری و رفتاری، بهبود دقت	Ensemble Learning and Image-Based Behavioral Data Transformation یادگیری مجموعه‌ای و تبدیل داده‌های رفتاری مبتنی بر تصویر	S. Zhao et al., 2025
Lack of complex models for non-linear data عدم استفاده از مدل‌های پیچیده برای داده‌های غیرخطی	Simple and executable, foundation for future studies ساده و قابل اجرا، پایه‌ای برای مطالعات بعدی	Educational Data Analysis with Classification Methods تحلیل داده‌های آموزشی با روش‌های دسته‌بندی	C. Jalota & R. Agrawal, 2019
Limited diversity in features and methods محدودیت در تنوع ویژگی‌ها و روش‌ها	Application in higher education, use of real-world datasets کاربرد در آموزش عالی، استفاده از مجموعه‌داده واقعی	Educational Data Mining for Student Success داده‌کاوی آموزشی برای موفقیت دانشجویان	N. Ndou et al., 2020
May not be suitable for large datasets ممکن است برای داده‌های بزرگ مناسب نباشد	Focus on predicting academic dropout, use of diverse ML techniques تمرکز روی پیش‌بینی افت تحصیلی، استفاده از تکنیک‌های متنوع ML	Predicting Student Dropout Status with EDA and ML پیش‌بینی وضعیت افت دانشجویان با EDA و ML	G. Jaiswal et al., 2020
Limited to existing features, lacks deep analysis محدود به ویژگی‌های موجود، بدون تحلیل عمیق	Focus on dropout, applicable in higher education institutions تمرکز روی dropout، قابل اعمال در مؤسسات آموزش عالی	Dropout Prediction with Data Mining پیش‌بینی ترک تحصیل با داده‌کاوی	W. W. Yaacob et al., 2020
Review study, no new model provided مطالعه مروری، بدون ارائه مدل جدید	Comprehensive analysis and comparison of methods تحلیل جامع روش‌ها و مقایسه آن‌ها	Review and Comparison of Educational Data Mining Techniques مرور و مقایسه تکنیک‌های داده‌کاوی آموزشی	Y. Zhang et al., 2021
Limited to exam data, potentially low generalizability محدود به داده‌های امتحانی، ممکن است تعمیم‌پذیری پایین	Use of diverse supervised methods استفاده از روش‌های نظارت‌شده متنوع	Exam Performance Prediction with Supervised Data Mining پیش‌بینی عملکرد امتحان با داده‌کاوی نظارت‌شده	K. M. Abiodun et al., 2021
Limited focus on imbalanced data, no hybrid model provided تمرکز محدود به داده نامتوازن، بدون ارائه مدل ترکیبی	Attention to the imbalance problem توجه به مشکل imbalance	Review of Methods for Handling Imbalanced Data in EDM بررسی روش‌های حل داده‌های نامتوازن در EDM	A. Al-Ashoor & S. Abdullah, 2022

Weaknesses نقاط ضعف	Strengths نقاط قوت	Proposed Method روش پیشنهادی	Author & Year نام نویسنده و سال
Requires specific data and processing, high complexity نیاز به داده و پردازش خاص، پیچیدگی بالا	Ability to model timing dynamics and neural activity توانایی مدل سازی دینامیک زمان بندی و فعالیت عصبی	Spiking Neural Network for Performance Prediction شبکه عصبی اسپایکینگ برای پیش بینی عملکرد	C. Liu et al., 2022
Review study, no empirical analysis مطالعه مروری، بدون تحلیل تجربی	Summarizing studies, identifying trends جمع بندی مطالعات، شناسایی روندها	Systematic Review of EDM (2015-2021) مرور نظام مند (۲۰۱۵-۲۰۲۱) EDM	M. H. bin Roslan & C. J. Chen, 2022
Limited focus, may not have generalizable results تمرکز محدود، ممکن است نتایج کلی نداشته باشد	Examining non-academic effects on performance بررسی اثرات غیر آموزشی بر عملکرد	Data Mining for the Effect of Social Media on Academic Performance داده کاوی برای اثر رسانه های اجتماعی بر عملکرد تحصیلی	S. Amjad et al., 2022
No advanced hybrid model بدون مدل ترکیبی پیشرفته	Standard and executable ML methods روش های ML استاندارد و قابل اجرا	Performance Prediction with ML Techniques پیش بینی عملکرد با تکنیک های ML	U. Verma et al., 2022
Requires image data, high complexity نیاز به داده های تصویری، پیچیدگی بالا	Ability to extract complex features توانایی استخراج ویژگی های پیچیده	Hybrid 2D CNN Model مدل هیبریدی ۲D CNN	S. Poudyal et al., 2022
Implementation complexity, requires fine-tuning پیچیدگی پیاده سازی، نیاز به تنظیم دقیق	Improved accuracy with ensemble learning, hybrid approach افزایش دقت با یادگیری ترکیبی، هیبرید	HELA: Hybrid Ensemble Algorithm الگوریتم هیبریدی مجموعه ای HELA	S. B. Keser & S. Aghalarova, 2022
Focus only on data balancing, no main model provided تمرکز فقط روی تعادل داده، بدون ارائه مدل اصلی	Examination of techniques to handle data imbalance بررسی تکنیک های مقابله با عدم تعادل داده	Comparison of Undersampling, Oversampling, and SMOTE مقایسه undersampling، oversampling و SMOTE	T. Wongvorachan et al., 2023
Limited to two subjects, potentially limited generalizability محدود به دو درس، ممکن است تعمیم پذیری محدود باشد	Practical for specific subjects, educational applicability کاربردی در درس های خاص، قابل استفاده آموزشی	Performance Prediction in English and Math with Data Mining پیش بینی عملکرد در انگلیسی و ریاضی با داده کاوی	M. H. B. Roslan & C. J. Chen, 2023

راهکارهای پیشنهادی

مدل ارائه شده در این مقاله یک روش یادگیری با نظارت گروهی است که در آن از ترکیب سه شبکه عصبی مختلف استفاده می شود. این مدل از سه مرحله پیش پردازش، انتخاب ویژگی، طبقه بندی و ترکیب روش ها و در نهایت ارزیابی تشکیل شده است. مراحل مدل ارائه شده در شکل ۱ ترسیم شده است.

نرمال سازی داده ها

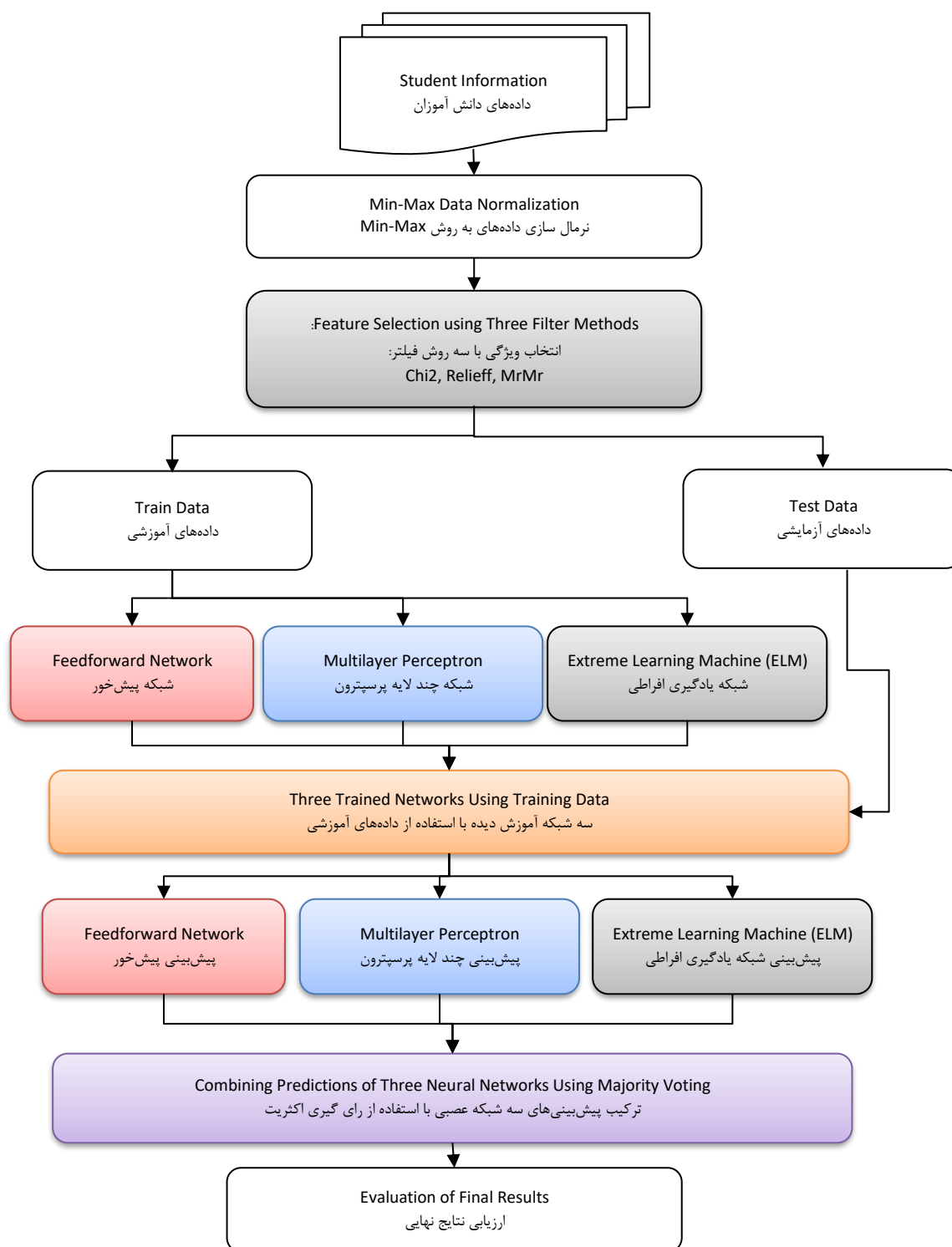
در ابتدا با استفاده از رابطه (۱) داده های دانش آموزان نرمال سازی می شوند.

$$X_{normal} = \frac{X_i - X_{min}}{X_{max} - X_{min}}, \quad i = 1, 2, \dots, N \quad (1)$$

در نرمال سازی بالا، X_i بیانگر مقدار هر داده در هر سطر از داده ها است. X_{normal} بیانگر مقدار نرمال شده داده ورودی است. X_{max} و X_{min} بیانگر حداقل و حداکثر داده ورودی است. این رابطه داده ها را به بازه [0,1] تبدیل می کند.

معرفی ویژگی ها و نمایش داده ها

مجموعه داده مورد استفاده در این پژوهش شامل مجموعه ای از ویژگی های تحصیلی، رفتاری و جمعیت شناختی است که جنبه های مختلف فرایند یادگیری و عملکرد تحصیلی دانشجویان را توصیف می کند. ویژگی های تحصیلی شامل نمرات پیشین، امتیازات ارزیابی دروس و شاخص های عملکرد تجمعی هستند که به طور مستقیم وضعیت تحصیلی دانشجو را بازنمایی می کنند. ویژگی های رفتاری الگوهای مشارکت آموزشی نظیر حضور در کلاس، میزان مشارکت و فعالیت های یادگیری را در بر می گیرند. همچنین، ویژگی های زمینه ای مانند سن، جنسیت و پیشینه آموزشی به منظور مدل سازی تفاوت های فردی لحاظ شده اند. در مجموع، مجموعه اولیه ویژگی ها شامل N ویژگی است که اگرچه توصیف جامعی از داده ها ارائه می دهد، اما می تواند شامل اطلاعات زائد یا هم بسته نیز باشد. از این رو، پیش از آموزش مدل های پیش بینی، استفاده از روش های انتخاب ویژگی به منظور شناسایی مؤثرترین ویژگی ها و کاهش افزونگی داده ها ضروری است.



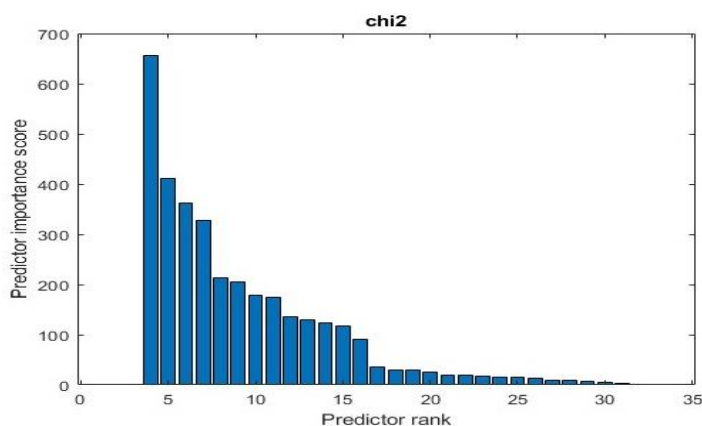
شکل ۱: دیاگرام مدل پیشنهادی

Fig. 1. Proposed Model Diagram

محاسبه کرده و بر اساس امتیاز محاسبه شده ویژگی‌ها را مرتب‌سازی می‌کنند و از ابتدای فهرست مرتب شده می‌توان ویژگی‌های برتر را انتخاب نمود. در این مطالعه از ۲۰ ویژگی استفاده شده است که تقریباً برابر با نیمی از ویژگی‌های داده‌های مورد استفاده است. در شکل‌های (۲) تا (۷) نمایه‌های ویژگی‌های برتر و نمای گرافیکی از وزن هر ویژگی‌ها به صورت مرتب شده توسط سه روش مختلف ارائه شده است.

انتخاب ویژگی به سه روش فیلتر

پس از نرمال‌سازی داده‌ها با استفاده از سه روش انتخاب ویژگی فیلتر، تعدادی از بهترین ویژگی‌ها انتخاب می‌شوند. هدف از انتخاب سه روش انتخاب ویژگی مختلف بررسی و مقایسه عملکرد آن‌ها با یکدیگر است. در الگوریتم Relief تعداد همسایگان برابر با ۱۰ در نظر گرفته شد. هر یک از این الگوریتم‌ها با استفاده از روابط خود امتیازی برای ویژگی‌ها



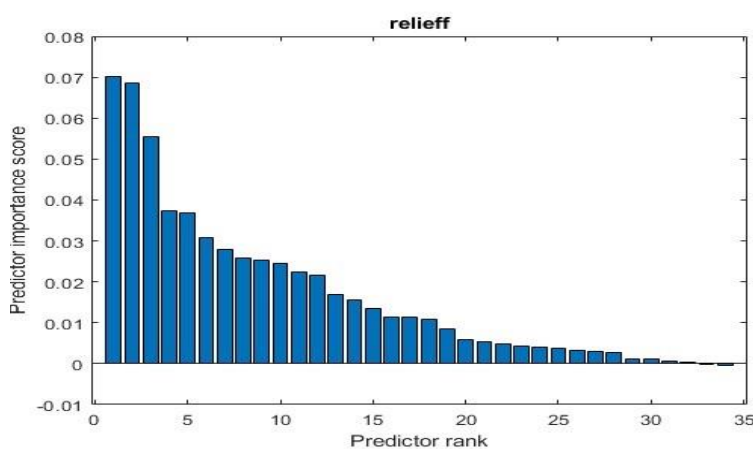
شکل ۳: امتیاز ویژگی‌ها تولید شده توسط Chi2

Fig. 5: Feature Scores Generated by Relief

23	29	30	24	15	28
	22	18	17	2	4
	21	14	27	16	6
	8	9	12	1	3
	31	10	25	5	11
	32	34	20	26	33
		7	19	13	

شکل ۲: فهرست ویژگی‌های برتر تولید شده توسط Chi2

Fig. 2: List of Top Features Generated by Chi2



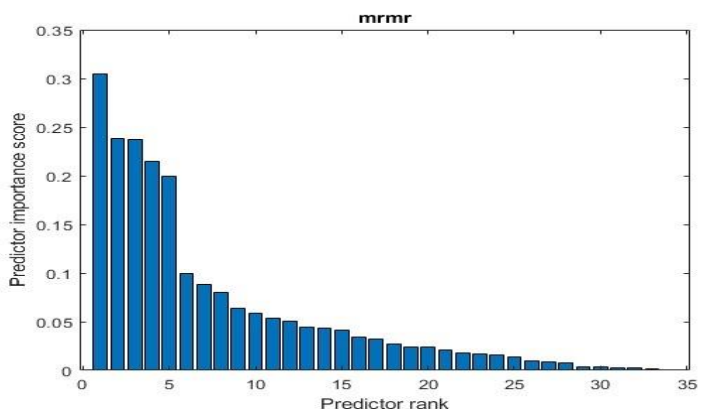
شکل ۵: امتیاز ویژگی‌ها تولید شده توسط Relief

Fig. 5: Feature Scores Generated by Relief

15	30	29	24	4	23
	17	32	12	33	34
	14	2	28	27	21
	8	9	22	10	3
	11	18	16	20	5
	26	6	31	25	13
		19	1	7	

شکل ۴: فهرست ویژگی‌های برتر تولید شده توسط Relief

Fig. 4: List of Top Features Generated by Relief



شکل ۷: امتیاز ویژگی‌ها تولید شده توسط mrmr

Fig. 7: Feature Scores Generated by mRMR

29	14	1	17	15	10
	23	16	18	30	25
	6	28	24	9	31
	12	22	11	7	3
	21	8	2	27	5
	13	4	33	20	26
		19	34	32	

شکل ۶: فهرست ویژگی‌های برتر تولید شده توسط mrmr

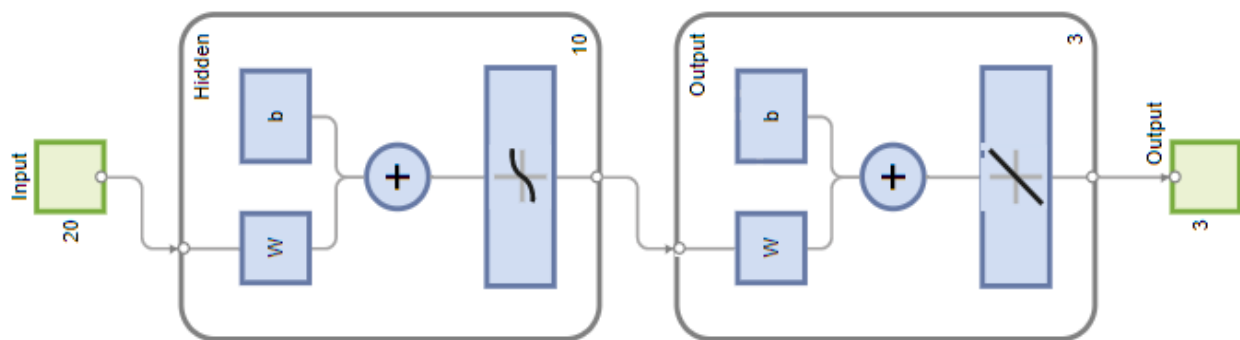
Fig. 6: List of Top Features Generated by mRMR

نورون‌های لایه ورودی برابر با ۲۰ است به این معنی که از کل ویژگی‌ها ۲۰ مورد انتخاب شده و به شبکه ارائه شده است و تعداد نورون‌های لایه پنهان ۱۰ است که به عنوان ورودی به شبکه داده می‌شود و نورون‌های لایه خروجی برابر با ۳ است که بیانگر سه کلاس است. در این دیاگرام w و b بیانگر وزن نورون‌های هر لایه و مقدار ثابت بایاس است. حداکثر تعداد گردش مرحله آموزش شبکه برابر با ۱۰۰۰ تعریف شده است.

طبقه‌بندی با سه شبکه عصبی

بعد از انتخاب ویژگی، از سه شبکه عصبی پیش‌خور را پیش‌بینی وضعیت تحصیلی دانش‌آموزان استفاده شده است.

- اولین شبکه پیش‌خور ساده است که دارای ۱ لایه ورودی، ۱ لایه خروجی و ۱ لایه پنهان است. در لایه پنهان از ۱۰ نورون استفاده شده است. در شکل (۸) نمودار این شبکه ارائه شده است. در این شکل تعداد



شکل ۸: نمودار شبکه پیش‌خور ساده با سه لایه
Fig. 8: Diagram of a Simple Three-Layer Feedforward Network

۱۰۰۰ آموزش شبکه انجام شده است. در شکل (۱۰) نمودار یادگیری شبکه ارائه شده است.

آخرین گام از مدل، ترکیب پیش‌بینی‌های انجام شده توسط مدل‌ها است که این بخش توسط رأی‌گیری اکثریت انجام می‌شود. در هر ردیف از داده‌های بخش آزمایش، کلاسی که دارای بیشترین رأی باشد به عنوان کلاس نهایی داده موردنظر انتخاب می‌شود. کلاس نهایی تولید شده توسط رأی‌گیری اکثریت خروجی نهایی مدل است. پس از تشکیل پیش‌بینی نهایی توسط رأی‌گیری اکثریت ارزیابی مدل انجام می‌شود.

ارزیابی و نتایج

برای ارزیابی عملکرد مدل ارائه شده از معیارهای دقت، صحت، فراخوان و میانگین f استفاده شده است که از طریق روابط (۴) تا (۷) محاسبه می‌شوند.

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

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

$$F_{score} = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (۷)$$

$$\text{Recall} = \frac{TP}{TP + FN} \quad (۶)$$

داده‌های این مقاله از نشانی Kaggle قابل دریافت است. مجموعه داده‌ها شامل داده‌های جمعیت‌شناختی، داده‌های اجتماعی - اقتصادی و کلان اقتصادی، داده‌های زمان ثبت نام دانشجو و داده‌های پایان ترم اول و دوم می‌باشد. منابع داده مورد استفاده شامل داده‌های داخلی و خارجی مؤسسه است و شامل داده‌هایی از (i) دستگاه مدیریت دانشگاهی (AMS)، (ii) دستگاه پشتیبانی از فعالیت‌های آموزشی، (iii) داده‌های سالانه اداره کل آموزش عالی (DGES) و (IV) پایگاه داده پرتغال در مورد داده‌های کلان اقتصادی است [۲۶].

داده‌ها به سوابق دانشجویان ثبت‌نام‌شده بین سال‌های تحصیلی ۲۰۰۸-۲۰۰۹ تا ۲۰۱۸-۲۰۱۹ اشاره دارد. این داده شامل داده‌های ۱۷ مدرک کارشناسی از رشته‌های مختلف دانش، مانند کشاورزی، طراحی، آموزش، پرستاری، روزنامه‌نگاری، مدیریت، خدمات اجتماعی و فناوری‌ها می‌شود. مجموعه داده نهایی یک فایل CSV در دسترس است و شامل ۴۴۲۴ رکورد با ۳۵ ویژگی است و حاوی مقادیر گم‌شده‌ای نیست [۲۶].

تابع فعال‌سازی لایه پنهان تابع تانژانت هایپربولیک است و تابع فعال‌سازی در لایه خروجی identity یا همان خطی است که به ترتیب در رابطه ۳(۲) و ۴(۴) ارائه شده است.

$$f(x) = \tanh(x) = \frac{2}{1 + e^{-2x}} - 1 \quad (۲)$$

$$f(x) = (x) \quad (۳)$$

- دومین شبکه، شبکه یادگیری ماشین افراطی است. در این شبکه تعداد نورون لایه مخفی برابر با ۵۰۰ تعریف شد. این شبکه مرحله یادگیری تکرارشونده ندارد و فرایند یادگیری آن تصادفی است.

- سومین شبکه، شبکه عصبی عمیق چندلایه پرسپترون است که لایه‌های آن در شکل (۹) ارائه شده است. این شبکه دارای ۷ لایه است و همچنین ۷/۹ هزار پارامتر یادگیری دارد. لایه‌های این شبکه به شرح زیر است:

○ لایه ورودی دارای ۲۰ نورون است و داده‌های ورودی انتخاب شده را دریافت می‌کند.

○ دومین یک‌لایه تمام متصل است که تعداد نورون‌های آن برابر با $32 \times \text{num_class}$ است. این لایه داده‌های ورودی را در وزن‌های مختص آن‌ها ضرب می‌کند و پس از جمع کردن مقادیر با مقدار بایاس آن‌ها را از یک تابع فعال‌سازی عبور می‌دهد.

○ سومین لایه نیز یک‌لایه تمام متصل دیگر است که تعداد نورون‌های آن برابر با $16 \times \text{num_class}$ است.

○ چهارمین لایه نیز یک‌لایه تمام متصل دیگر است که تعداد نورون‌های آن برابر با $8 \times \text{num_class}$ است.

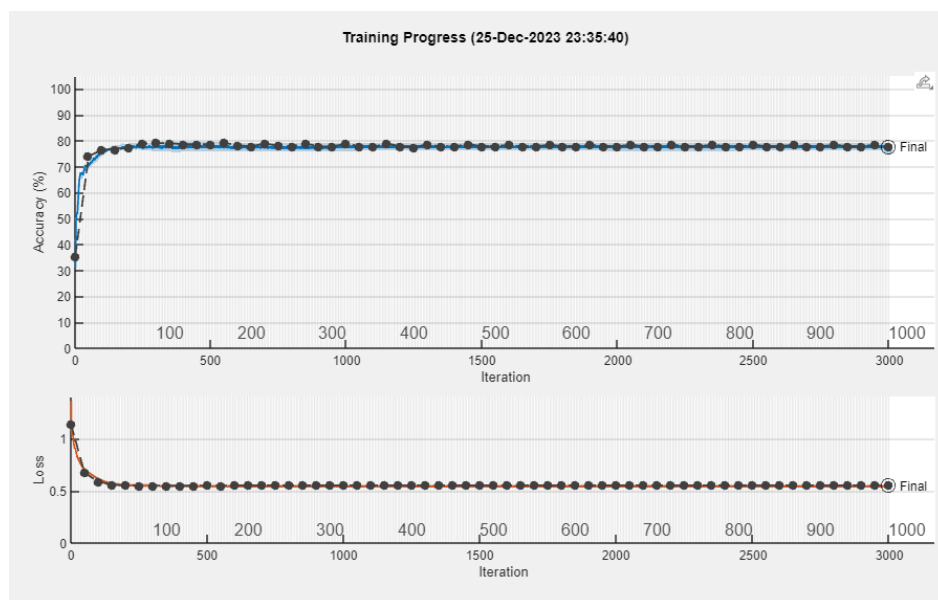
○ پنجمین لایه آخرین لایه تمام متصل است که تعداد نورون‌های ورودی آن برابر با تعداد کلاس مسئله است.

○ دولایه آخر در شبکه‌های طبقه‌بندی درج می‌شود و کلاس داده‌ها را محاسبه کرده و میزان خطای طبقه‌بندی شبکه در طبقه‌بندی داده‌ها را محاسبه می‌کند و به لایه‌های قبلی ارائه می‌دهد تا وزن‌های نورون‌های لایه‌های قبل به‌درستی تنظیم مجدد شوند. در لایه طبقه‌بندی تابع خطا Cross Entropy است وظیفه محاسبه خطای شبکه را دارد. پس از تعریف شبکه با استفاده از تابع Adam و داده‌های بخش آموزش و با تعداد Epoch برابر با ۱۰۰۰ و Minibatch برابر با



شکل ۹: لایه‌ها و پارامترهای شبکه عصبی عمیق چندلایه پرسپترون

Fig. 9: Layers and Parameters of the Deep Multilayer Perceptron Neural Network



شکل ۱۰: نمودار یادگیری شبکه MLP پیشنهادی

Fig. 10: Learning Curve of the Proposed MLP Network

حداقل افزونگی حداکثر همبستگی قرار می‌گیرد که به صحت ۷۴/۶۶٪ و معیار F برابر با ۶۶/۹۶٪ رسیده است. با مقایسه مقادیر سه روش انتخاب ویژگی فوق می‌توان مشاهده نمود که روش انتخاب ویژگی مناسب می‌تواند به‌خوبی منجر به افزایش عملکرد مدل شود. با مقایسه روش‌های فوق می‌توان گفت الگوریتم Relief بین ۳ تا ۶٪ منجر به بهبود مرحله انتخاب ویژگی و در نتیجه بهبود عملکرد کلی مدل شده است.

بررسی تعداد ویژگی منتخب

پس از بررسی و مقایسه روش‌های انتخاب ویژگی در آزمون بعد تعداد ویژگی بررسی شده است. نتایج این آزمایش در شکل (۱۲) گزارش شده است. این آزمایش فقط برای الگوریتم Relief انجام شد. در این آزمایش نیز نتایج رأی‌گیری اکثریت و ترکیب کلیه مدل‌ها گزارش شده است. نتایج این آزمایش در شکل (۱۲) نشان می‌دهد که در تعداد ۲۰ ویژگی بهترین نتایج حاصل شده است. بدترین نتایج متعلق به تعداد ۱۰ ویژگی است که صحت در این تعداد ویژگی برابر با ۷۳/۰۷٪ که نشان می‌دهد دانش کافی از داده‌ها حاصل نشده است. در تعداد ۳۵ ویژگی یعنی کل ویژگی‌های داده، صحت برابر با ۷۸.۲۸٪ و معیار F برابر با ۷۰/۱۳٪ شده است. در واقع با مقایسه مقدار حاصل شده با ۳۵ ویژگی می‌توان عملکرد مدل را قبل از انتخاب ویژگی و بعد از انتخاب ویژگی بررسی نمود. به این دلیل که تعداد ۳۵ ویژگی کل ویژگی‌های داده است در

در این مقاله از محیط متلب برای شبیه‌سازی روش‌ها استفاده شده است. کلیه آزمون‌های این فصل در یک دستگاه ۷ هسته با رم ۱۲ انجام شده است. برای آموزش شبکه‌های از ۹۰٪ داده‌ها برای آموزش و ۱۰٪ داده‌ها برای ارزیابی استفاده شده است. تعداد نورون برای شبکه ELM برابر با ۵۰۰، تعداد نورون در شبکه ANN برابر با ۱۰، تعداد نورون در شبکه MLP به ترتیب ۹۶-۴۸-۲۴-۳ تعریف گردید. تابع یادگیری شبکه MLP برابر با Adam در نظر گرفته شد. تعداد Epoch و سائز Minibatch برابر با ۱۰۰۰ تعریف شد. تعداد همسایگان برای انتخاب ویژگی Relief نیز برابر با ۱۰ تعریف شد.

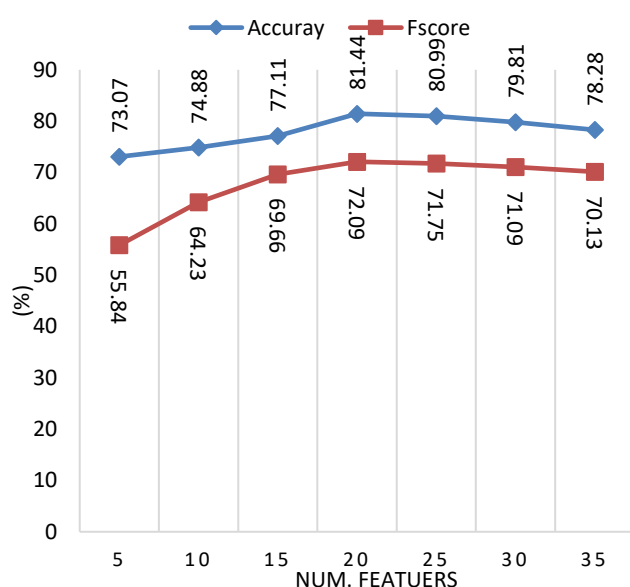
بررسی نوع الگوریتم انتخاب ویژگی

در یکی از مراحل مدل پیشنهادی، از سه روش انتخاب ویژگی فیلتر استفاده شده است تا مشخص شود کدام یک از این روش‌ها می‌توانند بهتر از دیگری عمل کنند. نتایج این آزمایش در شکل (۱۱) ارائه شده است. در این آزمایش نتایج رأی‌گیری اکثریت گزارش شده است.

نتایج این نمودار نشان می‌دهد که در بین سه روش انتخاب ویژگی الگوریتم Relief دارای نتایج بهتری در مقایسه با روش‌های دیگر است. با استفاده از این الگوریتم و انتخاب ۲۰ ویژگی برتر مدل به صحت ۸۱/۴۴٪ و به معیار F برابر با ۷۲/۰۹٪ دست‌یافته است. در رتبه دوم الگوریتم کای ۲ قرار می‌گیرد که این الگوریتم به ترتیب به صحت ۷۷/۸۲٪ و معیار F برابر با ۶۷/۵۳٪ رسیده است. در رتبه سوم الگوریتم

نتایج متعلق به ELM است. در بین این سه شبکه ماشین یادگیری افراطی توانسته است به صحت ۸۲/۸٪ دست یابد؛ یعنی بین ۲٪ تا ۴٪ توانسته است صحت را در مقایسه با دو شبکه عصبی دیگر این مطالعه افزایش دهد.

در معیار دقت نیز شبکه ماشین یادگیری افراطی به دقت ۸۱/۴۸٪ رسیده است که در مقایسه با دو شبکه دیگر بین ۷٪ تا ۹٪ توانسته است دقت را بهبود دهد. در واقع نتایج بالا نشان می‌دهد انتخاب معماری مناسب و شبکه عصبی مناسب در نتایج مدل‌ها مؤثر است و انتخاب معماری درست می‌تواند باعث افزایش عملکرد مدل به طور چشمگیری شود. در شکل‌های (۱۴) تا (۱۶) ماتریس درهم‌ریختگی این سه شبکه ارائه شده است.



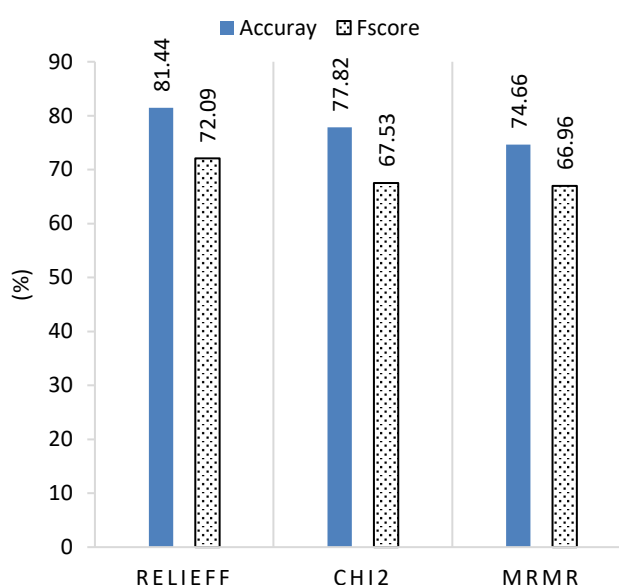
شکل ۱۲: تعداد ویژگی مناسب در مرحله انتخاب ویژگی
Fig. 12: Optimal Number of Features in the Feature Selection Stage

نتیجه می‌توان گفت با تعداد ۳۵ ویژگی یعنی مرحله انتخاب ویژگی نادیده گرفته شده است. با مقایسه تعداد ۳۵ و ۲۰ ویژگی می‌توان گفت که بعد از انتخاب ویژگی صحت مدل ۳/۱۶٪ و معیار F مدل نیز ۱/۹۶٪ افزایش داشته است. در نتیجه مرحله انتخاب ویژگی منجر به موفقیت بیشتر مدل در طبقه‌بندی داده‌ها شده است.

بررسی نوع شبکه عصبی عمیق

در این مقاله از سه نوع شبکه عصبی استفاده شده است که هر یک عملکرد خود را دارد. در شکل (۱۳) نتایج این آزمایش نیز ارائه شده است.

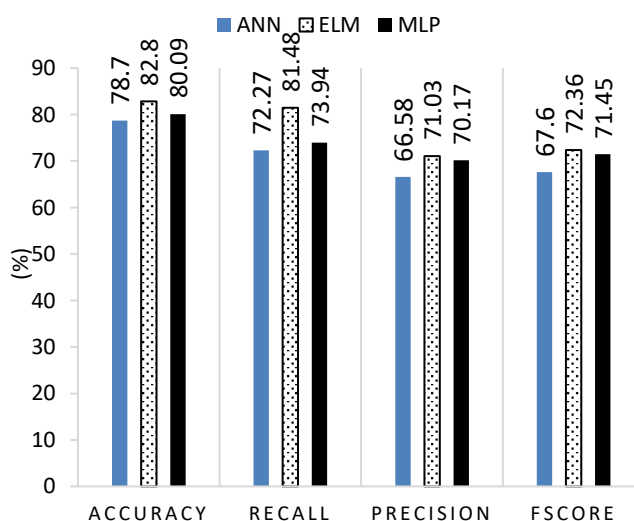
نتایج سه شبکه عصبی نشان می‌دهد که در بین شبکه عمیق MLP و شبکه پیش‌خور ANN و شبکه ماشین یادگیری افراطی ELM بهترین



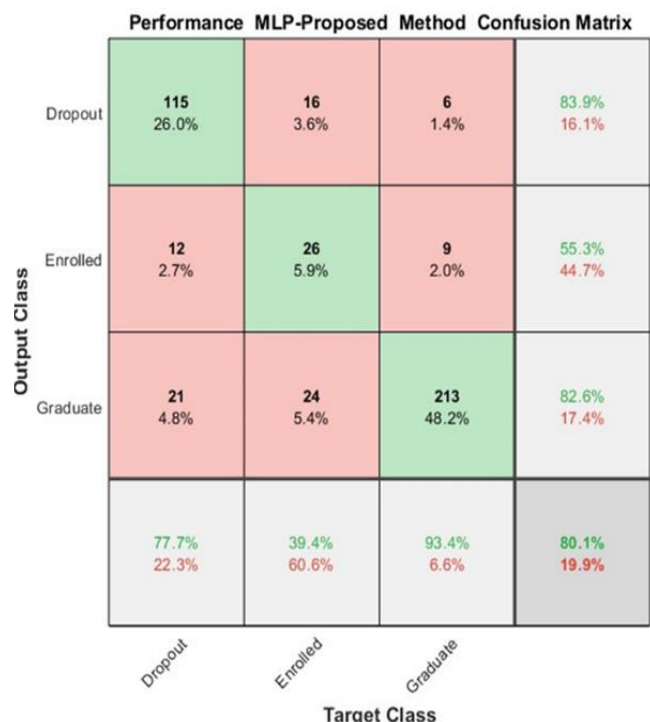
شکل ۱۱: بررسی نوع الگوریتم انتخاب ویژگی فیلتر
Fig. 11: Evaluation of Filter-Based Feature Selection Algorithm Type

Performance ELM-Proposed Method Confusion Matrix				
Output Class	1	2	3	
	131 29.6%	25 5.7%	11 2.5%	78.4% 21.6%
	3 0.7%	20 4.5%	2 0.5%	80.0% 20.0%
	14 3.2%	21 4.8%	215 48.6%	86.0% 14.0%
Target Class				
	88.5% 11.5%	30.3% 69.7%	94.3% 5.7%	82.8% 17.2%

شکل ۱۴: ماتریس درهم‌ریختگی شبکه ELM
Fig. 14: Confusion Matrix of the ELM Network



شکل ۱۳: مقایسه بین سه شبکه عصبی پیشنهادی
Fig. 13: Comparison Among the Three Proposed Neural Networks



شکل ۱۶: ماتریس درهم‌ریختگی شبکه MLP
Fig. 16: Confusion Matrix of the MLP Network

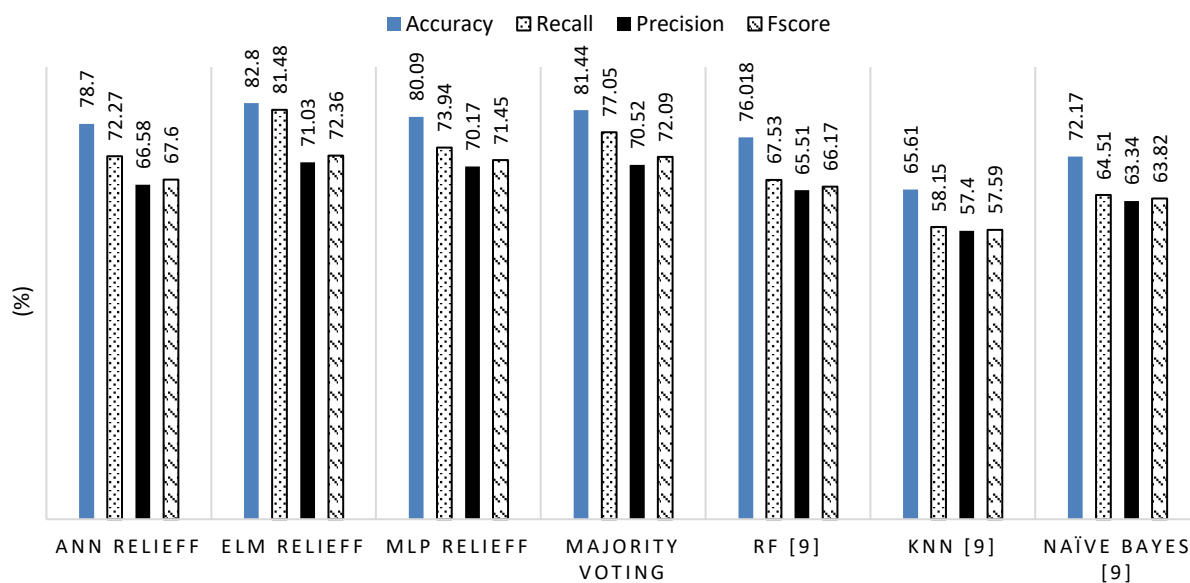


شکل ۱۵: ماتریس درهم‌ریختگی شبکه ANN
Fig. 15: Confusion Matrix of the ANN Network

مقایسه نتایج

اکثریت قرار می‌گیرد که در مقایسه با دیگر روش‌ها دارای نتایج بهتری است. علت موفقیت مدل وابسته به دو مرحله انتخاب ویژگی مناسب و طبقه‌بندی مناسب است. استفاده از سه شبکه عصبی مختلف و ترکیب نتایج آن‌ها با یکدیگر باعث شده مدل توانایی مناسبی در پیش‌بینی وضعیت تحصیلی دانش‌آموزان داشته باشد. خلاصه نتایج فوق در ادامه بیان شده است:

در آخرین بخش مقایسه بین روش‌های این مقاله و مطالعه [۷] انجام شده است. نتایج این آزمایش در شکل (۱۷) ارائه شده است. مقایسه نتایج شکل (۱۷) نشان می‌دهد که در بین همه روش‌ها، ماشین یادگیری افراطی دارای نتایج بهتری است. این شبکه منجر به بهبود و افزایش روش رأی‌گیری اکثریت شده است. در رتبه دوم رأی‌گیری



شکل ۱۷: مقایسه نتایج روش‌ها با یکدیگر
Fig. 17: Comparison of the Methods' Results

است. همچنین نتایج نشان داد که شبکه عصبی یادگیری افراطی در مقایسه با دو شبکه دیگر نتایج بسیار امیدوارکننده‌ای ارائه می‌دهد. این شبکه به صحت ۸۲/۸٪ رسیده است. در پایان نیز مقایسه روش‌ها با روش‌های دیگر ثابت کرد شبکه‌های عصبی در مقایسه با مدل‌های یادگیری ماشین سنتی دارای عملکرد بسیار بهتری هستند. باتوجه به محدودیت‌های پژوهش حاضر، در کارهای آینده پیشنهاد می‌شود مدل‌های پیشنهادی بر روی دیتاست‌های واقعی و متنوع آموزشی ارزیابی شوند تا تعمیم‌پذیری نتایج افزایش یابد. همچنین، امکان بررسی و ترکیب الگوریتم‌های پیشرفته یادگیری ماشین و شبکه‌های عصبی عمیق، از جمله Transformers، برای بهبود دقت پیش‌بینی وجود دارد. افزودن ویژگی‌های رفتاری و تعاملات دانشجویان در محیط‌های آموزشی آنلاین می‌تواند به غنی‌تر شدن مدل‌ها کمک کند و امکان ارائه توصیه‌های عملی به دانشجویان و مؤسسات آموزشی را فراهم نماید. علاوه بر این، تحلیل تفسیرپذیری مدل‌ها و ارزیابی آنها با معیارهایی مانند سرعت پردازش و کارایی هزینه - فایده، می‌تواند به کاربرد عملی مدل‌ها در سیستم‌های آموزشی کمک شایانی کند.

مشارکت نویسندگان

راهنمایی در اجرای طرح تحقیق، تحلیل داده‌ها و نگارش مقاله توسط دکتر حسن ضیافت ارائه شده است. جمع‌آوری داده‌ها توسط آقای دکتر قاسم آذری آرانی انجام شده است.

تشکر و قدردانی

نویسندگان مایل‌اند قدردانی خود را از همه همکاران، داوران و دوستانی که نظرات و پیشنهادهای ارزشمندی ارائه دادند و کیفیت این مقاله را بهبود بخشیدند، ابراز کنند.

تعارض منافع

نویسندگان مایل‌اند اعلام کنند که هیچ تضاد منافعی در ارتباط با انتشار این مقاله وجود ندارد. نویسندگان تأیید می‌کنند که پژوهش حاضر به‌صورت عینی و بدون هیچ‌گونه تأثیر از روابط مالی یا شخصی انجام شده است که ممکن است به‌عنوان تضاد منافع تلقی شود.

منابع و مأخذ

- [1] Shorfuzzaman M, Hossain MS, Nazir A, Muhammad G, Alamri A. Harnessing the power of big data analytics in the cloud to support learning analytics in mobile learning environment. *Comput Human Behav.* 2019;92:578–88. doi: 10.1016/j.chb.2018.07.002.
- [2] Capuano N, Toti D. Experimentation of a smart learning system for law based on knowledge discovery and cognitive computing. *Comput Human Behav.* 2019;92:459–67. doi: 10.1016/j.chb.2018.03.034.

شبکه عصبی ANN در مقایسه با سه روش یادگیری ماشین سنتی در مطالعه یا گچی و همکاران [۷] به طور میانگین صحت را ۷/۴۳٪، دقت را ۸/۸۷٪، فراخوانی را ۴/۴۹٪ و معیار F را ۵/۰۷٪ افزایش داده است. شبکه عصبی یادگیری افراطی در مقایسه با سه روش یادگیری ماشین سنتی در مطالعه یا گچی و همکاران [۹] به طور میانگین صحت را ۱۱/۵۳٪، دقت را ۱۸/۰۸٪، فراخوانی را ۸/۹۴٪ و معیار F را ۹/۸۳٪ افزایش داده است. شبکه عصبی عمیق MLP در مقایسه با سه روش یادگیری ماشین سنتی در مطالعه یا گچی و همکاران [۹] به طور میانگین صحت را ۸/۸۲٪، دقت را ۱۰/۵۴٪، فراخوانی را ۸/۰۸٪ و معیار F را ۸/۹۲٪ افزایش داده است. ترکیب سه شبکه عصبی با روش رأی‌گیری اکثریت در مقایسه با سه روش یادگیری ماشین سنتی در مطالعه یا گچی و همکاران [۹] به طور میانگین صحت را ۱۰/۱۷٪، دقت را ۱۳/۶۵٪، فراخوانی را ۸/۴۳٪ و معیار F را ۹/۵۶٪ افزایش داده است. نتایج شکل (۱۷) نشان می‌دهد که شبکه‌های عصبی در مقایسه با روش‌های یادگیری ماشین عملکرد بسیار بهتری دارند. به طور میانگین می‌توان با مقایسه کلیه روش‌ها گفت که صحت بیش از ۷٪ افزایش داشته است و دقت نیز بالای ۴٪ افزایش یافته است.

نتیجه‌گیری

پیش‌بینی وضعیت تحصیلی دانش‌آموزان اهمیت زیادی در سیاق سیاست‌های آموزشی و تربیتی دارد. این اطلاعات به مدیران آموزشی، معلمان، و والدین کمک می‌کند تا با دستیابی به درک عمیق‌تر از علل موفقیت یا شکست تحصیلی، برنامه‌ها و سیاست‌های بهینه‌تری را پیاده‌سازی کنند. از طرف دیگر، این اطلاعات به محققان و سیاست‌گذاران اجتماعی نیز کمک می‌کند تا درک بهتری از چگونگی تأثیر عوامل اجتماعی، اقتصادی، فردی، و فرایندهای آموزشی بر تحصیلات دانش‌آموزان کسب کنند. یادگیری ماشین و شبکه‌های عصبی به‌عنوان ابزارهای قدرتمند در تجزیه و تحلیل داده‌های حجیم به‌شدت به چشم می‌آیند. این تکنیک‌ها از توانایی‌های الگوریتم‌های پیچیده و تشخیص الگوهای پنهان در داده‌ها برای پیش‌بینی وضعیت‌های آتی استفاده می‌کنند. به‌عنوان مثال، با تحلیل داده‌های مربوط به حضور و غیاب، نمرات، رفتارهای تحصیلی، و عوامل دیگر، این سیستم‌ها می‌توانند به مدارس و سازمان‌های آموزشی کمک کنند تا دانش‌آموزان با خطر احتمالی از دست رفتن وضعیت تحصیلی را شناسایی کنند و برنامه‌های پیشگیرانه مناسبی اجرا کنند. برای رسیدن به این هدف در این مقاله یک روش ترکیبی مبتنی به سه شبکه عصبی ماشین یادگیری افراطی، شبکه عصبی پیش‌خور ساده و شبکه عصبی عمیق چندلایه پرسپترون استفاده گردید. در این مدل نتایج روش‌ها با استفاده از رأی‌گیری اکثریت ترکیب شد. همچنین یک مرحله انتخاب ویژگی مبتنی بر سه روش مختلف در مدل درج شد تا مشخص شود کدام الگوریتم دارای نتایج بهتری است. نتایج بررسی‌های مختلف نشان داد که در بین سه روش انتخاب ویژگی، الگوریتم Relief دارای نتایج بهتری

using data mining techniques. J Phys Conf Ser. 2020; 1496: 012005. doi: 10.1088/1742-6596/1496/1/012005.

[15] Zhang Y, Yun Y, An R, Cui J, Dai H, Shang X. Educational data mining techniques for student performance prediction: method review and comparison analysis. Front Psychol. 2021; 12: 698490. doi: 10.3389/fpsyg.2021.698490.

[16] Abiodun KM, Adeniyi EA, Aremu DR, Awotunde JB, Ogbuji E. Predicting students performance in examination using supervised data mining techniques. In: Misra S, Oluranti J, Damaševičius R, Maskeliūnas R, editors. Informatics and intelligent applications: proceedings of the First International Conference, ICIIA 2021. Cham: Springer; 2022. p. 63–77. doi: 10.1007/978-3-030-95630-1_5.

[17] Al-Ashoor A, Abdullah S. Examining techniques to solving imbalanced datasets in educational data mining systems. Int J Comput. 2022;21(2):205–13. doi: 10.47839/ijc.21.2.2589.

[18] Liu C, Wang H, Du Y, Yuan Z. A predictive model for student achievement using spiking neural networks based on educational data. Appl Sci. 2022;12(8):3841. doi: 10.3390/app12083841.

[19] Bin Roslan MH, Chen CJ. Educational data mining for student performance prediction: a systematic literature review (2015–2021). Int J Emerg Technol Learn. 2022;17(5):147–79. doi: 10.3991/ijet.v17i05.27685.

[20] Amjad S, Younas M, Anwar M, Shaheen Q, Shiraz M, Gani A. Data mining techniques to analyze the impact of social media on academic performance of high school students. Wirel Commun Mob Comput. 2022;2022:9299115. doi: 10.1155/2022/9299115.

[21] Verma U, Garg C, Bhushan M, Samant P, Kumar A, Negi A. Prediction of students' academic performance using machine learning techniques. In: Proceedings of the 2022 International Mobile and Embedded Technology Conference (MECON); 2022 Mar 10–11; Noida, India. IEEE; 2022. p. 151–6. doi: 10.1109/MECON53876.2022.9751956.

[22] Poudyal S, Mohammadi-Aragh MJ, Ball JE. Prediction of student academic performance using a hybrid 2D CNN model. Electronics. 2022;11(7):1005. doi: 10.3390/electronics11071005.

[23] Keser SB, Aghalarova S. HELA: a novel hybrid ensemble learning algorithm for predicting academic performance of students. Educ Inf Technol. 2022;27:4521–52. doi: 10.1007/s10639-021-10780-0.

[24] Wongvorachan T, He S, Bulut O. A comparison of undersampling, oversampling, and SMOTE methods for dealing with imbalanced classification in educational data mining. Information. 2023;14(1):54. doi: 10.3390/info14010054.

[25] Bin Roslan MH, Chen CJ. Predicting students' performance in English and Mathematics using data mining techniques. Educ Inf Technol. 2023;28(2):1427–53. doi: 10.1007/s10639-022-11259-2.

[3] Viberg O, Hatakka M, Bälter O, Mavroudi A. The current landscape of learning analytics in higher education. Comput Human Behav. 2018;89:98–110. doi: 10.1016/j.chb.2018.07.027.

[4] Waheed H, Hassan SU, Aljohani NR, Hardman J, Alelyani S, Nawaz R. Predicting academic performance of students from VLE big data using deep learning models. Comput Human Behav. 2020;104:106189. doi: 10.1016/j.chb.2019.106189.

[5] Bernacki ML, Chavez MM, Uesbeck PM. Predicting achievement and providing support before STEM majors begin to fail. Comput Educ. 2020;158:103999. doi: 10.1016/j.compedu.2020.103999.

[6] Batool S, Rashid J, Nisar MW, Kim J, Kwon HY, Hussain A. Educational data mining to predict students' academic performance: a survey study. Educ Inf Technol. 2023;28(1):905–71. doi: 10.1007/s10639-022-11152-y.

[7] Yağcı M. Educational data mining: prediction of students' academic performance using machine learning algorithms. Smart Learn Environ. 2022;9:11. doi: 10.1186/s40561-022-00192-z.

[8] Zhou H, Shen S, Su Y, Miao Y, Liu Q, Zhu L, et al. LLM-EPSP: large language model empowered early prediction of student performance. Inf Process Manag. 2026;63(1):104351. doi: 10.1016/j.ipm.2025.104351.

[9] Eriyandi V, Ahmad AN. A comparative machine learning and deep learning models with ElasticNet regularization for predicting student outcomes: LGBM, CatBoost, ANN, DNN, and WDN. Int J Adv Comput Inform. 2026;2(2):96–107. doi: 10.71129/ijaci.v2i2.pp96-107.

[10] Zhao S, Zhou D, Wang H, Chen D, Yu L. Enhancing student academic success prediction through ensemble learning and image-based behavioral data transformation. Appl Sci. 2025;15(3):1231. doi: 10.3390/app15031231.

[11] Jalota C, Agrawal R. Analysis of educational data mining using classification. In: Proceedings of the 2019 International Conference on Machine Learning, Big Data, Cloud and Parallel Computing (COMITCon); 2019 Feb 14–16; Faridabad, India. IEEE; 2019. p. 243–7. doi: 10.1109/COMITCon.2019.8862214.

[12] Ndou N, Ajoodha R, Jadhav A. Educational data-mining to determine student success at higher education institutions. In: Proceedings of the 2020 2nd International Multidisciplinary Information Technology and Engineering Conference (IMITEC); 2020 Nov 25–27; Kimberley, South Africa. IEEE; 2020. p. 1–8. doi: 10.1109/IMITEC50163.2020.9334139.

[13] Jaiswal G, Sharma A, Sarup R. Machine learning in higher education: predicting student attrition status using educational data mining. In: Solanki A, Kumar S, Nayyar A, editors. Handbook of research on emerging trends and applications of machine learning. Hershey (PA): IGI Global; 2020. p. 27–46. doi: 10.4018/978-1-5225-9643-1.ch002.

[14] Wan Yaacob WF, Mohd Sobri N, Nasir SAM, Norshahidi ND, Wan Husin WZ. Predicting student drop-out in higher institution

Azariaran, Gh. Assistant Professor, Computer Engineering, Technical and Vocational University, (TVU), Tehran, Iran
 ✉ ghazariarani@tvu.ac.ir



حسن ضیافت دارای مدرک دکتری تخصصی در رشته مهندسی کامپیوتر از دانشگاه کاشان است. ایشان بیش از ۲۰ سال سابقه فعالیت آموزشی و پژوهشی در دانشگاه‌ها و مراکز علمی را دارند و با مجلات علمی مختلف نیز همکاری داشته‌اند. از سال ۱۴۰۰ تاکنون به‌عنوان عضو

هیئت علمی در گروه کامپیوتر دانشگاه آزاد اسلامی واحد نطنز مشغول به تدریس و پژوهش هستند. زمینه‌های تخصصی ایشان شامل داده‌کاوی، محاسبات ابری، یادگیری ماشین و هوش مصنوعی است.

Ziafat, H. Assistant Professor, Computer Engineering, Islamic Azad University, Natanz, Iran
 ✉ dr.ziafat@iau.ac.ir

[26] Realinho V, Machado J, Baptista L, Martins MV. Predicting student dropout and academic success. Data. 2022;7(11):146. doi: 10.3390/data7110146.

معرفی نویسندگان

AUTHOR(S) BIOSKETCHES



قاسم آذری آرانی دارای مدرک دکتری تخصصی فناوری اطلاعات از دانشگاه آزاد قم است. ۳۲ سال فعالیت آموزشی و پژوهشی و همکاری با مجلات علمی دارد و از سال ۱۳۹۷ تا اکنون به‌عنوان استادیار در گروه مهندسی برق و کامپیوتر و فناوری اطلاعات دانشگاه ملی مهارت دانشکده فنی شهید رجایی کاشان مشغول به فعالیت هستند. زمینه‌های تخصصی ایشان عبارت‌اند از: هوش مصنوعی، یادگیری ماشین و رباتیک، تبدیل متن به گفتار، مدیریت ارتباط با مشتریان، مدیریت زنجیره تأمین، سیستم‌های عامل.

Citation (Vancouver): Ziafat H1, Azariarani Gh. [Improving the accuracy of predicting students' academic success and failure based on the efficiency of deep feedforward neural networks]. *Tech. Edu. J.* 2026; 20(2): 271-286

<https://doi.org/10.22061/tej.2026.12612.3316>

