

Hypertension Prediction in Primary School Students Using an Ensemble Machine Learning Method

Besharati Reza¹, Tahmasbi HamidReza^{2*}

• Received: 13 Sep 2022

• Accepted: 1 Nov 2022

Introduction: The prevalence of hypertension in children is increasing, and this complication is considered the most important risk factor for cardiovascular diseases in older age. Early detection and control of hypertension can prevent its progress and reduce its consequences. Machine learning methods can help predict this complication promptly and reduce cost and time. This study aimed to provide a model based on ensemble machine learning methods to more accurately predict the hypertension of primary school children.

Method: This is an applied developmental study that was conducted using the information of 1287 primary school children aged 7-13 years in Kashmar city. After data preprocessing, to achieve a more accurate diagnosis of hypertension in children, the output results of five common machine learning methods in disease diagnosis including decision tree, naive Bayesian, nearest neighbors, artificial neural network, and support vector machine using weighted majority voting method were combined.

Results: The results showed that the accuracy, sensitivity, and specificity of the proposed model were 90.31%, 80.65%, and 93.54%, respectively, and compared to similar studies it performed better.

Conclusion: The proposed model can better predict and diagnose hypertension in children and improve accuracy and reduce the error rate. This model can be a useful and early tool in the diagnosis of hypertension in children, reducing the consequences and costs of this complication and being a big step in the fight against hypertension.

Keywords: Hypertension, Primary School Students, Machine Learning Methods, Prediction

• **Citation:** Besharati R, Tahmasbi HR. Hypertension Prediction in Primary School Students Using an Ensemble Machine Learning Method. *Journal of Health and Biomedical Informatics* 2022; 9(3): 148-57. :doi 10.34172/jhbmi.2022.05. [In Persian]

1. PhD in Health Care Management, Assistant Professor, Department of Nursing, Kashmar Branch, Islamic Azad University, Kashmar, Iran

2. PhD in Computer Engineering, Assistant Professor, Department of Computer Engineering, Kashmar Branch, Islamic Azad University, Kashmar, Iran

*Corresponding Author: HamidReza Tahmasbi

Address: Kashmar Branch, Islamic Azad University, Kashmar, Iran

• Tel: 09151046117

• Email: htahma@gmail.com

پیش‌بینی فشار خون بالا در کودکان دبستانی با استفاده از ترکیب روش‌های یادگیری ماشین

رضا بشارتی^۱، حمیدرضا طهماسبی^{۲*}

• پذیرش مقاله: ۱۴۰۱/۸/۱۰

• دریافت مقاله: ۱۴۰۱/۶/۲۲

مقدمه: شیوع فشار خون بالا در کودکان رو به افزایش است و این عارضه مهم‌ترین عامل خطر برای بیماری‌های قلبی-عروقی در سنین بالاتر به شمار می‌رود. تشخیص به موقع فشار خون بالا و کنترل آن می‌تواند جلوی پیشرفت آن را گرفته و پیامدهای ناشی از آن را کاهش دهد. روش‌های یادگیری ماشین می‌توانند به پیش‌بینی به موقع این عارضه کمک کرده و باعث کاهش هزینه و زمان گردند. این مطالعه با هدف ارائه مدلی مبتنی بر ترکیب روش‌های یادگیری ماشین برای تشخیص و پیش‌بینی دقیق‌تر فشار خون کودکان دبستانی انجام شد.

روش: این مطالعه از نوع کاربردی-توسعه‌ای بوده که با استفاده از اطلاعات ۱۲۸۷ نفر از کودکان دبستانی ۷ تا ۱۳ ساله شهر کاشمر انجام شده است. پس از پیش پردازش داده‌ها، برای تشخیص دقیق‌تر کودکان مبتلا به فشار خون بالا نتایج خروجی پنج روش یادگیری ماشین متداول در تشخیص بیماری‌ها، شامل درخت تصمیم، بیزین ساده، نزدیکترین همسایه‌ها، شبکه عصبی مصنوعی و ماشین بردار پشتیبان با استفاده از روش رأی‌گیری اکثریت وزن‌دار ترکیب می‌شوند.

نتایج: نتایج نشان داد که دقت (Accuracy)، حساسیت (Sensitivity) و ویژگی (Specificity) در مدل پیشنهادی به ترتیب ۹۰/۳۱، ۸۰/۶۵ و ۹۳/۵۴ درصد بوده و در مقایسه با مطالعات مشابه، عملکرد بهتری دارد.

نتیجه‌گیری: مدل پیشنهادی بهتر می‌تواند پیش‌بینی و تشخیص فشار خون بالا در کودکان را انجام داده و باعث بهبود دقت و کاهش میزان اشتباه گردد. این مدل می‌تواند به عنوان یک ابزار مفید و زود هنگام در تشخیص فشار خون بالا در کودکان، از پیامدها و هزینه‌های ناشی از این عارضه بکاهد و گام بزرگی در مبارزه با فشار خون بالا باشد.

کلیدواژه‌ها: فشار خون بالا، کودکان دبستانی، روش‌های یادگیری ماشین، پیش‌بینی

• **ارجاع:** بشارتی رضا، طهماسبی حمیدرضا. پیش‌بینی فشار خون بالا در کودکان دبستانی با استفاده از ترکیب روش‌های یادگیری ماشین. مجله انفورماتیک سلامت و زیست پزشکی

doi 10.34172/jhbmi.2022.05: ۱۴۸-۵۷ (۳)۹؛ ۱۴۰۱

۱. دکتر مدیریت خدمات بهداشتی و درمانی، استادیار، گروه پرستاری، واحد کاشمر، دانشگاه آزاد اسلامی، کاشمر، ایران

۲. دکتر مهندسی کامپیوتر، استادیار، گروه مهندسی کامپیوتر، واحد کاشمر، دانشگاه آزاد اسلامی، کاشمر، ایران

* **نویسنده مسئول:** حمیدرضا طهماسبی

آدرس: کاشمر، دانشگاه آزاد اسلامی کاشمر

• **Email:** htahma@gmail.com

• **شماره تماس:** ۰۹۱۵۱۰۴۶۱۱۷

مقدمه

فشار خون بالا یک بیماری مزمن است که در سراسر جهان شایع بوده و هر سال از سطح سنی مبتلایان به آن کاسته می‌شود [۱]. این عارضه نقش مؤثری در افزایش خطر حمله قلبی، نارسایی کلیوی، بیماری‌های عروق کرونری قلب و مرگ و میر دارد [۲]. [۳] و یک دغدغه اصلی بهداشت جهانی به شمار می‌رود [۴]. با توجه به گزارش سازمان بهداشت جهانی حدود ۱۳ درصد از مرگ و میرهای جهان ناشی از فشارخون بالا می‌باشد [۵، ۳]. اکثر افراد مبتلا به فشار خون بالا هیچ علامت و نشانه‌ای ندارند [۶] و در نتیجه این عارضه یک قاتل خاموش محسوب می‌شود [۳]. تشخیص زودهنگام فشار خون بالا و نظارت و کنترل آن می‌تواند جلوی پیشرفت آن را گرفته و پیامدهای ناشی از آن را کاهش دهد و به همین دلیل مورد توجه بسیاری از پژوهشگران قرار گرفته است [۲]. مطالعات بالینی نشان می‌دهند برای افرادی که در مرحله پیش بالینی فشارخون بالا هستند یا در معرض خطر ابتلا به فشار خون هستند، پیشرفت بیماری ممکن است به طور قابل توجهی از طریق تغییر سبک زندگی یا با یک درمان دارویی مؤثر کاهش یابد [۴].

امروزه فشار خون بالا دیگر یک بیماری بزرگسالان نیست و تعداد فزاینده‌ای از کودکان و نوجوانان امروزی به دلایلی از قبیل تغذیه، سبک زندگی و عدم فعالیت فیزیکی دچار آن می‌شوند. در مطالعه جامعی که درباره شیوع افزایش فشارخون بالا در بین کودکان و نوجوانان ایرانی انجام شده، زنگ خطر افزایش سالانه ۰/۲۱ درصدی شیوع فشار خون بالا در کودکان ایرانی را به صدا در آورده است. متأسفانه فشار خون بالا در کودکان تا زمانی که به حدی نرسد که زندگی آنان را به خطر بیندازد و یا تا زمانی که به بزرگسالی نرسند، تشخیص داده نمی‌شود. با توجه به پیامدهای درازمدت فشارخون بالای کنترل نشده برای سلامتی و همچنین این واقعیت که فشارخون بالا در کودکان و نوجوانان یک سیگنال تشخیصی برای بسیاری از بیماری‌های مهم پزشکی است، تشخیص زودهنگام و صحیح آن ضروری می‌باشد [۷، ۸].

در چند سال اخیر روش‌های دسته‌بندی در یادگیری ماشین به‌عنوان یک راهکار مهم در تشخیص و کنترل بیماری‌ها مورد توجه قرار گرفته‌اند. یادگیری ماشین یکی از شاخه‌های هوش مصنوعی می‌باشد و نشان داده شده است که استفاده از آن در توسعه ابزارها و مدل‌های تشخیص و پیش‌بینی بیماری‌ها نسبت به روش‌های دستی و مرسوم، عملکرد بهتری دارد [۹]. به طوری که سازمان بهداشت جهانی به کاربرد آن در حوزه پزشکی و

سودمندی دانش استخراج شده از داده‌های پزشکی و سلامت در تشخیص و پیش‌بینی بیماری‌ها تأکید کرده است [۱۰]. مطالعات اخیر نشان داده‌اند که روش‌های یادگیری ماشین برای پیش‌بینی، تشخیص و کنترل فشار خون بالا نیز مؤثر و مفید بوده و مورد توجه پژوهشگران زیادی قرار گرفته‌اند. به عنوان مثال Oanh و همکاران [۶] از روش‌های یادگیری بیزن ساده (Naive NB(Bayes)، شبکه‌های عصبی مصنوعی (Artificial ANN(Neural Networks)، ماشین بردار پشتیبان (SVM(Support Vector Machine)، نزدیک‌ترین همسایه (KNN(K-Nearest Neighbors)، درخت تصمیم (DT(Decision Tree)، رگرسیون خطی (Logistic LR(Regression)، جنگل تصادفی (Random Forest RF، رأی‌گیری و (Extreme Gradient Boosting XGBoost برای پیش‌بینی فشار خون بالا در مراجعین به بیمارستانی در ویتنام استفاده کردند. یافته‌های آن‌ها بیانگر سودمندی این روش‌های یادگیری در پیش‌بینی فشار خون بالا می‌باشد. آن‌ها نشان دادند که روش جنگل تصادفی نسبت به سایر روش‌های مقایسه شده عملکرد بهتری داشته است. Islam و همکاران [۹] از روش‌های درخت تصمیم، جنگل تصادفی، XGBoost، تقویت گرادیان (GBM)، رگرسیون خطی (LR) و روش آنالیز تشخیص خطی (LDA) برای پیش‌بینی فشار خون بالا در مجموعه‌ی داده‌های بیماران جنوب آسیا و همچنین شناسایی عوامل پرخطر مؤثر در بروز فشارخون بالا استفاده کردند. آن‌ها نشان دادند که روش‌های XGBoost، GBM، رگرسیون خطی و LDA با دقتی حدود ۹۰ درصد از دو روش درخت تصمیم و جنگل تصادفی بهتر بوده‌اند. این روش‌ها ویژگی‌های سن و شاخص توده بدنی (BMI) را به عنوان عوامل مؤثر در بروز فشار خون بالا در مجموعه داده بررسی شده شناسایی کردند. Zhao و همکاران [۱۱] مدلی برای پیش‌بینی فشارخون بالا با استفاده از روش یادگیری جنگل تصادفی از روی ویژگی‌هایی که به راحتی از افراد و بدون نیاز به ابزار تخصصی خاصی قابل دست‌یافتنی است، پیشنهاد کردند. این مدل دقتی حدود ۸۲ درصد داشته و ویژگی‌های BMI، سن، دور کمر و سابقه خانوادگی فشار خون را به عنوان عوامل اصلی بروز فشار خون بالا در مجموعه داده مورد نظر شناسایی می‌کند. آن‌ها در مطالعه خود نشان دادند که پیش‌بینی فشار خون بالا بدون نیاز به دسترسی به داده‌های بالینی و ژنتیکی، با استفاده از الگوریتم یادگیری جنگل تصادفی امکان‌پذیر بوده و از عملکرد خوبی برخوردار است.

SVM و DT روی مجموعه‌ی داده‌ها اعمال و سپس خروجی آن‌ها توسط روش رگرسیون لجستیک برای به دست آوردن نتیجه نهایی پیش‌بینی، استفاده می‌شود. دقت این روش بر روی دو مجموعه داده برای تشخیص فشارخون بالا به ترتیب ۸۵/۷۳ درصد و ۷۵/۷۸ درصد بوده است.

با توجه به ضرورت تشخیص به موقع فشار خون بالا در کودکان و اهمیت دقت و اعتماد در این تشخیص، هدف این مطالعه، ارائه مدلی مبتنی بر ترکیب روش‌های یادگیری ماشین برای پیش‌بینی و دسته‌بندی دقیق‌تر فشار خون کودکان دبستانی ۷ تا ۱۳ ساله می‌باشد. مدل‌های ارائه شده در مطالعات قبلی افراد را به دو دسته‌ی فشار خون طبیعی و فشار خون بالا دسته‌بندی می‌کنند. در حالی که در مدل پیشنهادی این مطالعه، عمل تشخیص و پیش‌بینی به صورت جزئی‌تر انجام می‌شود. به طوری که این مدل افراد را با توجه به دستورالعمل آکادمی پزشکی اطفال آمریکا [۱۷] به چهار دسته فشار خون طبیعی، پیش فشار خون بالا، فشار خون بالای مرحله ۱ و فشار خون بالای مرحله ۲ دسته‌بندی می‌کند.

روش

این پژوهش، یک مطالعه کاربردی-توسعه‌ای است که با پیشنهاد مدلی مبتنی بر ترکیب روش‌های یادگیری ماشین به دسته‌بندی و پیش‌بینی فشار خون بالا در کودکان دبستانی می‌پردازد. نمونه آماری تعداد ۱۲۸۷ نفر از دانش‌آموزان پسر و دختر ۷ تا ۱۳ ساله دبستان‌های شهر کاشمر می‌باشند. داده‌ها پس از کسب مجوز از اداره کل آموزش و پرورش استان خراسان رضوی و اداره آموزش و پرورش شهرستان کاشمر و با اخذ رضایت نامه والدین آن‌ها جمع‌آوری شد. ویژگی‌های جمع‌آوری شده مربوط به دانش‌آموزان در جدول ۱ مشاهده می‌شود. مقادیر ویژگی‌های سن، جنسیت، سابقه مصرف داروهای مؤثر بر فشارخون، میزان مصرف نمک، سابقه‌ی بیماری‌های دیابت، قلبی، کلیوی و سایر بیماری‌ها، و سابقه خانوادگی پرفشاری خون از طریق پرسشنامه‌ای که در اختیار والدین دانش‌آموزان قرار گرفت، به دست آمد. ملاک محاسبه سن دانش‌آموزان، سن شناسنامه‌ای آن‌ها بود. قد هر دانش‌آموز در حالت کاملاً ایستاده بدون کفش و کلاه، به نحوی که پاها کنار هم قرار گرفته و فرد کاملاً به دیوار چسبیده و نگاه رو به جلو و مستقیم دارد، با قدسنج سکا (Seca) و با دقت نیم سانتی‌متر اندازه‌گیری شد. وزن دانش‌آموزان بر حسب کیلوگرم و با حداقل لباس ممکن، بدون کفش و با ترازوی دیجیتالی سکای آلمانی اندازه‌گیری شد. نبض

روش یادگیری شبکه عصبی مصنوعی رایج‌ترین روش یادگیری ماشین است که برای تشخیص و پیش‌بینی فشار خون بالا استفاده شده است [۲]. دهقاندار و همکاران [۱۲] از روش یادگیری شبکه عصبی مصنوعی برای تشخیص چاقی و فشار خون بالا در دانش‌آموزان اصفهانی بین ۷ تا ۱۸ سال استفاده کردند. ارزیابی آن‌ها نشان داد که مدل آن‌ها می‌تواند فشار خون سیستمولیک و دیاستولیک را به ترتیب با دقت ۷۴ درصد و ۷۹ درصد تشخیص دهد. Chai و همکاران [۸] از روش‌های مختلف یادگیری ماشین برای تشخیص فشار خون بالا در نوجوانان ۱۳ تا ۱۷ ساله استفاده کردند. مقایسه‌ها و ارزیابی‌های آنان نشان داد که روش یادگیری ماشین LightGBM (روش تقویت گرادیان سبک) از عملکرد بهتری در مقایسه با سایر روش‌ها برای تشخیص فشار خون بالا برخوردار بوده است. این پژوهشگران در مطالعه دیگری [۱۳] با توسعه روش یادگیری ماشین شبکه‌های عصبی مصنوعی، مدلی با دقت حدود ۷۶ درصد برای پیش‌بینی فشار خون بالا در نوجوانان ۱۳ تا ۱۷ ساله ارائه نمودند.

در تعدادی از مدل‌های موفق پیشنهاد شده برای پیش‌بینی فشارخون بالا از ترکیب چندین روش یادگیری مختلف استفاده شده است. رویکرد ترکیبی یکی از روش‌های یادگیری ماشین است که نتیجه نهایی پیش‌بینی و تشخیص بر اساس ترکیب نتایج خروجی چند روش یادگیری حاصل می‌شود. این ترکیب یک راهکار قدرتمند برای افزایش دقت دسته‌بندی است [۱۴]. مطالعات انجام شده نشان می‌دهند که این رویکرد در مقایسه با هر یک از روش‌های یادگیری استفاده شده در ترکیب، دقت بالاتری داشته [۱۵] و تأثیر قابل توجهی در بهبود تشخیص فشار خون بالا و عوامل مؤثر بر آن دارد [۱۶]. Fang و همکاران [۴] مدلی ترکیبی پیشنهاد کردند که با ترکیب نتایج خروجی دو روش یادگیری KNN و LightGBM فشار خون بالا را پیش‌بینی می‌کند. ارزیابی آن‌ها نشان داد که این مدل پیشنهادی دقتی حدود ۸۶ درصد دارد. در این مدل، هر یک از روش‌های KNN و LightGBM با احتمالی پیش‌بینی می‌کنند که نمونه ورودی متعلق به کدام دسته می‌باشد. نتیجه نهایی پیش‌بینی بر اساس میانگین وزنی این احتمالات مشخص می‌شود. این مدل از ویژگی سن و شاخص‌های خونی افراد، بدون در نظر گرفتن ویژگی‌هایی از قبیل سابقه خانوادگی و یا سبک زندگی که به طور مستقیم با فشارخون بالا مرتبط هستند، برای پیش‌بینی فشار خون بالا استفاده می‌کند. Fitriyani و همکاران [۱۶] روش ترکیبی ارائه کردند که ابتدا الگوریتم‌های یادگیری MLP،

نیز با استفاده از کورنومتر برای یک دقیقه کامل اندازه‌گیری شد. شاخص توده بدنی (BMI) برای هر دانش‌آموز به صورت نسبت وزن بر حسب کیلوگرم به مجذور قد بر حسب متر مربع محاسبه گردید. فشار خون سیستولیک و دیاستولیک بر حسب میلی‌متر جیوه (mmHg)، در وضعیت نشسته و از بازوی راست در حالی که بازو در سطح قلب قرار داشت و پس از ایجاد آرامش روحی و جسمی در دانش‌آموز و استراحت ۵ دقیقه‌ای، ۲ بار و به فاصله ۵ دقیقه با فشارسنج عقربه‌ای ژاپنی و گوشی پزشکی ژاپنی طوری اندازه‌گیری شد که کاف فشارسنج تقریباً ۴۰ درصد عرض بازو و طول آن ۸۰ تا ۱۰۰ درصد محیط بازو را پوشانده و در بالای حفره آنتی کوبیتال بسته شود. کاف به اندازه ۳۰ تا ۴۰ میلی‌متر جیوه بالای فشار خون سیستولیک، باد شده و با سرعت حدود ۳ میلی‌متر جیوه در ثانیه کاهش داده می‌شد، صدای اول کورتکوف به عنوان فشار سیستولیک و صدای چهارم یا پنجم کورتکوف به عنوان فشار دیاستولیک ثبت شد. میانگین دو نوبت

اندازه‌گیری، به عنوان فشارخون نهایی فرد منظور شد. از مقادیر فشارخون سیستولیک و دیاستولیک برای تعیین ویژگی تشخیص یعنی دسته‌بندی دانش‌آموزان توسط فرد خبره به یکی از چهار دسته فشار خون طبیعی، پیش فشار خون بالا، فشار خون بالای مرحله ۱ و فشارخون بالای مرحله ۲ بر اساس جدول ۲ برگرفته از دستورالعمل آکادمی پزشکی اطفال آمریکا [۱۷] استفاده شد و کاربرد دیگری در این مطالعه نداشته‌اند. بر این اساس، تعداد ۱۰۶۱ نفر وضعیت فشار خون طبیعی، ۶۲ نفر پیش فشارخون بالا، ۱۲۱ نفر فشار خون بالای مرحله ۱ و ۴۳ نفر فشارخون بالای مرحله ۲ داشته‌اند. مقادیر ۱۶ ویژگی ردیف ۱ تا ۱۶ در جدول ۱ به عنوان ورودی مدل پیشنهادی در نظر گرفته شد و خروجی مدل، نتیجه تشخیص (فشارخون طبیعی/پیش فشارخون بالا/فشار خون بالای مرحله ۱/فشارخون بالای مرحله ۲) می‌باشد.

جدول ۱: ویژگی‌های جمع‌آوری شده مربوط به داده‌های مورد بررسی

ردیف	عنوان ویژگی	بازدهی مقادیر	توضیحات
۱	سن	عدد صحیح بین ۷ تا ۱۳	ورودی مدل
۲	جنسیت	۰: پسر، ۱: دختر	ورودی مدل
۳	قد	عدد صحیح بر حسب سانتی‌متر	ورودی مدل
۴	وزن	عددی با یک رقم اعشار بر حسب کیلوگرم	ورودی مدل
۵	نبض	عدد صحیح	ورودی مدل
۶	شاخص توده بدنی (BMI)	عددی با دو رقم اعشار بر حسب کیلوگرم بر مترمربع	ورودی مدل
۷	سابقه مصرف داروهای مؤثر بر فشارخون	۰: خیر، ۱: بلی	ورودی مدل
۸	میزان مصرف نمک	۰: کم نمک، ۱: معمولی، ۲: پر نمک	ورودی مدل
۹	سابقه بیماری دیابت	۰: خیر، ۱: بلی	ورودی مدل
۱۰	سابقه بیماری قلبی	۰: خیر، ۱: بلی	ورودی مدل
۱۱	سابقه بیماری کلیوی	۰: خیر، ۱: بلی	ورودی مدل
۱۲	سابقه سایر بیماری‌ها	۰: خیر، ۱: بلی	ورودی مدل
۱۳	سابقه پر فشاری والدین	۰: هیچ‌کدام، ۱: پدر یا مادر، ۲: هر دو	ورودی مدل
۱۴	سابقه پر فشاری برادر/خواهر	۰: خیر، ۱: بلی	ورودی مدل
۱۵	سابقه پر فشاری پدر بزرگ/مادر بزرگ	۰: خیر، ۱: بلی	ورودی مدل
۱۶	سابقه پر فشاری عمو/دایی/عمه/خاله	۰: خیر، ۱: بلی	ورودی مدل
۱۷	فشار خون سیستولیک	عدد صحیح بر حسب میلی‌متر جیوه (mmHg)	استفاده شده برای دسته‌بندی نمونه‌ها توسط فرد خبره
۱۸	فشار خون دیاستولیک	عدد صحیح بر حسب میلی‌متر جیوه (mmHg)	استفاده شده برای دسته‌بندی نمونه‌ها توسط فرد خبره
۱۹	تشخیص	۰: فشارخون طبیعی (۱۰۶۱ نمونه) ۱: پیش فشارخون بالا (۶۲ نمونه) ۲: فشارخون بالای مرحله ۱ (۱۲۱ نمونه) ۳: فشارخون بالای مرحله ۲ (۴۳ نمونه)	خروجی مدل

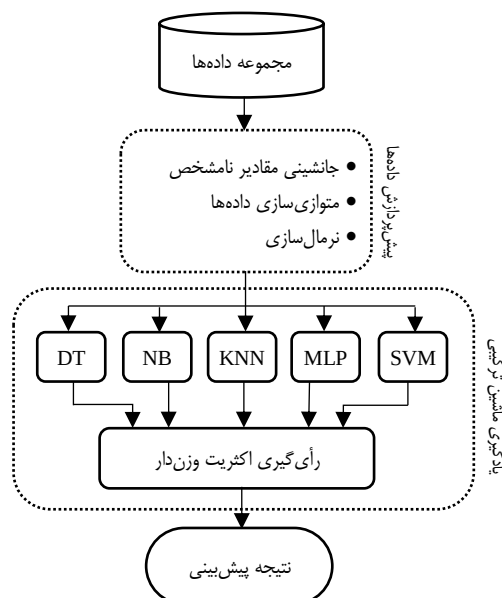
جدول ۲: دسته‌بندی فشارخون دانش‌آموزان براساس معیارهای فشار سیستولیک و دیاستولیک

دسته فشار خون	فشار خون سیستولیک (mmHg)	فشار خون دیاستولیک (mmHg)
فشار خون طبیعی	کمتر از ۱۲۰	و کمتر از ۸۰
پیش فشارخون بالا	۱۲۰ تا ۱۳۹	و کمتر از ۸۰
فشارخون بالای مرحله ۱	۱۳۰ تا ۱۳۹	یا ۸۰ تا ۸۹
فشارخون بالای مرحله ۲	۱۴۰ به بالا	یا ۹۰ به بالا

شکل ۱ ساختار مدل پیشنهادی را نشان می‌دهد. ابتدا پیش پردازش داده‌ها به عنوان یکی از مهمترین اقدامات در فرایند داده‌کاوی و یادگیری ماشین بر روی مجموعه داده‌ها انجام شد. بدین منظور مقدار ویژگی میزان مصرف نمک در ۵۸ نمونه از دانش‌آموزان نامشخص بود که برای هر نمونه با مقداری که بین نمونه‌های هم دسته با آن نمونه، بیشترین تکرار را داشته است، جایگزین شد. در مجموعه داده جمع‌آوری شده، تعداد نمونه‌های از دسته دارای فشارخون طبیعی در مقایسه با نمونه‌های سه دسته دیگر به طور چشم‌گیری بیشتر می‌باشد. این عدم توازن در داده‌ها که در اغلب داده‌های پزشکی وجود دارد، می‌تواند دقت دسته‌بندی مدل را تحت تأثیر قرار دهد و معمولاً نمونه‌های متعلق به دسته‌های کوچک به عنوان نمونه‌های دسته با بیشترین تعداد نمونه دسته‌بندی می‌شوند [۱۸، ۱۶]. بر این اساس، این خطر در مجموعه داده‌های جمع‌آوری شده وجود دارد که مدل، نمونه‌های متعلق به دسته‌های پیش فشار خون بالا، فشار خون بالای مرحله ۱ و فشار خون بالای مرحله ۲ را اشتباهاً به عنوان دسته فشار خون طبیعی، دسته‌بندی کند.

برای مقابله با این مشکل، داده‌ها با استفاده از تکنیک

متوازن‌سازی معروف (Synthetic Minority Over-sampling Technique) SMOTE [۱۹] متوازن شد. مطالعات انجام شده نشان می‌دهند که تکنیک SMOTE می‌تواند دقت الگوریتم‌های دسته‌بندی را برای دسته اقلیت بهبود بخشد [۲۱، ۲۰، ۱۸]. در این تکنیک برای دسته با تعداد نمونه‌های کمتر، نمونه‌های جدیدی در همسایگی نمونه‌های موجود در آن دسته تولید می‌شود. در نتیجه تعداد نمونه‌های متعلق به هر یک از دسته‌ها، متوازن می‌شوند. با این کار، تعداد نمونه‌های دسته‌های پیش فشار خون بالا، فشار خون مرحله ۱ و فشار خون مرحله ۲ به ترتیب به ۱۰۵۴، ۱۰۶۴ و ۱۰۳۲ نمونه افزایش یافت. پس از متوازن‌سازی داده‌ها، با استفاده از روش نرمال‌سازی min-max، مقادیر ویژگی‌های ورودی مدل (ویژگی‌های ردیف ۱ تا ۱۶ در جدول ۱) نرمال شدند. به طوری که همه مقادیر در بازه [۰، ۱] قرار می‌گیرند. در این روش نرمال‌سازی، کمترین مقدار یک ویژگی، تبدیل به صفر و بیشترین مقدار ویژگی تبدیل به یک می‌شود و سایر مقادیر آن ویژگی به مقادیر بین صفر و یک تبدیل می‌شوند.



شکل ۱: ساختار مدل پیشنهادی

است، اضافه می‌شود. در پایان، دسته‌ای که بیشترین وزن را داشته باشد به عنوان نتیجه نهایی دسته‌بندی فشارخون انتخاب می‌شود. از ماژول‌های نرم‌افزار داده‌کاوی متن باز Weka نسخه ۳,۷,۸ با همان پارامترهای پیش‌فرض و همچنین زبان برنامه‌نویسی پایتون (Python 3.10.8) برای پیاده‌سازی و آزمایش‌ها استفاده شد. برای ارزیابی مدل پیشنهادی، دو سوم از مجموعه داده‌ها به صورت تصادفی به عنوان مجموعه داده آموزشی و یک سوم بقیه به عنوان مجموعه آزمایش انتخاب شدند. برای به دست آوردن پارامترهای بهینه و وزن هر روش یادگیری از روش اعتبارسنجی متقابل ۱۰ تکه برابر (10-fold cross-validation) [۲۳] استفاده شد. در بین مقادیر مختلف بررسی شده برای تعداد k همسایه در روش یادگیری KNN، مقدار بهینه $k=10$ انتخاب شد.

کارایی با استفاده از معیارهای ارزیابی دقت (Accuracy)، حساسیت (Sensitivity) و ویژگی (Specificity) [۲۴] مورد بررسی قرار گرفت. از آنجایی که تعداد دسته‌ها، بیش از دو دسته می‌باشند، میانگین مقادیر هر معیار در دسته‌ها بر اساس روش macro-average [۲۵] به صورت زیر محاسبه شد.

در مرحله بعد، فرآیند یادگیری و ساخت مدل انجام شد. در این مرحله، ابتدا پنج روش یادگیری ماشین قدرتمند و متداول در تشخیص و دسته‌بندی بیماری‌ها [۶] شامل درخت تصمیم C4.5 (DT (Decision Tree)، بیزین ساده (Naive K-Nearest)، نزدیک‌ترین k همسایه (NB (Bayesian KNN (Neighbors)، شبکه‌های عصبی مصنوعی چندلایه MLP (Multi-Layer perceptron) و ماشین بردار پشتیبان (SVM (Support Vector Machine) بر روی مجموعه‌ی داده‌ها اعمال و دسته‌بندی نمونه‌ها را انجام می‌دهند. به هر روش یادگیری با توجه به میزان دقتش در دسته‌بندی، یک وزن تخصیص می‌یابد. به طوری که مجموع وزن همه پنج روش برابر یک می‌باشد. سپس خروجی روش‌های یادگیری با استفاده از روش رأی‌گیری اکثریت وزن‌دار با هم ترکیب شده و بر اساس آن نتیجه‌ی نهایی دسته‌بندی و تشخیص مشخص می‌شود. روش رأی‌گیری اکثریت وزن‌دار معروف‌ترین و پرکاربردترین روش ترکیب الگوریتم‌های یادگیری ماشین می‌باشد [۲۲، ۱۴]. بر مبنای این روش در مدل پیشنهادی، ابتدا وزن صفر به هر یک از چهار دسته داده می‌شود. سپس وزن هر روش یادگیرنده به وزن دسته‌ای که توسط آن روش به عنوان خروجی تعیین شده

$$\text{Accuracy} = \frac{\sum_{i=1}^m \frac{TP_i + TN_i}{TP_i + FN_i + FP_i + TN_i}}{m} \quad (1)$$

$$\text{Sensitivity}_{\text{macro}} = \frac{\sum_{i=1}^m \frac{TP_i}{TP_i + FN_i}}{m} \quad (2)$$

$$\text{Specificity}_{\text{macro}} = \frac{\sum_{i=1}^m \frac{TN_i}{TN_i + FP_i}}{m} \quad (3)$$

است که به اشتباه در دسته دیگری دسته‌بندی شده‌اند.

نتایج

جدول ۳ نتیجه عملکرد مدل پیشنهادی در مقایسه با هر یک از روش‌های یادگیری استفاده شده در ترکیب را بر حسب معیارهای دقت، حساسیت و ویژگی به دست آمده نشان می‌دهد. مقادیر پررنگ در جدول بیانگر بیشترین مقدار هستند.

که m به معنی تعداد دسته‌ها می‌باشد. TP_i (True Positive) به معنی تعداد نمونه‌های متعلق به دسته i است که درست دسته‌بندی شده‌اند. FP_i (False Positive) به معنی تعداد نمونه‌هایی است که به اشتباه به عنوان دسته i دسته‌بندی شده‌اند. TN_i (True Negative) به معنی تعداد نمونه‌های متعلق به سایر دسته‌ها است که در دسته i دسته‌بندی نشده‌اند. FN_i (Negative) به معنی تعداد نمونه‌های متعلق به دسته i

جدول ۳: کارایی مدل پیشنهادی و روش‌های یادگیری شرکت کننده در ترکیب (درصد)

روش	دقت	حساسیت	ویژگی
درخت تصمیم (DT)	۸۴/۴۱	۶۸/۸۳	۸۹/۶۱
بیزین ساده (NB)	۶۷/۱۵	۳۴/۶۹	۷۸/۲۲
نزدیک‌ترین k همسایه (KNN)	۷۶/۶۳	۵۳/۳۲	۸۴/۴۵
شبکه عصبی پرسپترون چند لایه (MLP)	۸۴/۹۲	۶۹/۸۹	۸۹/۹۵
ماشین بردار پشتیبان (SVM)	۷۵/۷۱	۵۱/۴۴	۸۳/۸۳
مدل پیشنهادی	۹۰/۳۱	۸۰/۶۵	۹۳/۵۴

جدول ۴ مقایسه کارایی مدل پیشنهادی با چند روش جدید دیگر برای تشخیص و پیش‌بینی فشارخون بالا را نشان می‌دهد. این روش‌های مقایسه شده بر خلاف مدل پیشنهادی، نمونه‌ها را به دو دسته طبیعی و فشارخون بالا دسته‌بندی می‌کنند. به عبارت

دیگر، این روش‌ها نمونه‌های از نوع پیش فشارخون بالا، فشارخون مرحله ۱ و فشارخون مرحله ۲ را به عنوان فشارخون بالا در نظر می‌گیرند.

جدول ۴: کارایی مدل پیشنهادی در مقایسه با سایر روش‌ها (درصد)

روش	دقت	حساسیت	ویژگی
Zhao و همکاران [۱۱]	۸۷/۱۱	۷۴/۲۲	۹۱/۴
Chai و همکاران [۱۳]	۸۵/۱۵	۷۰/۵	۹۰/۱۱
Fang و همکاران [۴]	۸۹/۴۴	۷۸/۸۸	۹۲/۹۶
Fitriyani و همکاران [۱۶]	۸۷/۶۶	۷۵/۳۴	۹۱/۷۸
مدل پیشنهادی	۹۰/۳۱	۸۰/۶۵	۹۳/۵۴

بحث و نتیجه‌گیری

در این مطالعه، با توجه به ضرورت تشخیص و پیش‌بینی به موقع فشارخون بالا در کودکان دبستانی، یک مدلی ترکیبی ارائه شد که بتواند فشارخون در کودکان را در یکی از چهار دسته فشارخون طبیعی، پیش فشارخون بالا، فشارخون بالای مرحله ۱ و فشارخون بالای مرحله ۲ تشخیص دهد. به منظور بهبود دقت در مدل پیشنهادی، نتایج خروجی پنج روش یادگیری ماشین متداول و قوی در تشخیص بیماری‌ها شامل DT، NB، KNN، MLP و SVM با استفاده از روش رأی‌گیری اکثریت وزن‌دار ترکیب می‌شوند. مقایسه مقادیر دقت، حساسیت و ویژگی هر یک از این پنج روش یادگیری در جدول ۳ نشان می‌دهد که روش MLP از عملکرد بهتری برخوردار بوده است و پس از آن روش DT کارایی خوبی داشته است. این یافته‌ها، عملکرد مؤثر شبکه‌های عصبی مصنوعی MLP را در تشخیص و پیش‌بینی فشارخون بالا تأیید می‌کنند. هر چند مقادیر معیارهای مقایسه شده در روش MLP اختلاف چندانی با روش DT ندارد، ولی مطالعات انجام شده نشان می‌دهند که حتی یک بهبود جزئی در مقادیر این معیارها در کاربردهای حیاتی از قبیل پزشکی درخور

توجه می‌باشد. در بین پنج روش یادگیری شرکت کننده در ترکیب، روش NB با دقت ۶۷/۱۵ درصد، حساسیت ۳۴/۶۹ درصد و ویژگی ۷۸/۲۲ درصد، کمترین کارایی را داشته است. از طرفی مقایسه‌ی مدل پیشنهادی با هر یک از پنج روش یادگیری در جدول ۳ نشان می‌دهد که مدل پیشنهادی دقت، حساسیت و ویژگی بالاتری داشته و با اختلاف چشمگیری نسبت به سایر روش‌ها از عملکرد مطلوب‌تری برخوردار بوده است. برتری مدل پیشنهادی در مقایسه با هر یک از روش‌های یادگیری شرکت کننده در ساخت مدل، صحت این ادعا که ترکیب روش‌های یادگیری ماشین می‌توانند بر محدودیت‌های هر یک از این روش‌ها غلبه کرده و باعث بهبود دقت در تشخیص و پیش‌بینی بیماری‌ها گردد را ثابت می‌کند.

بررسی نتایج به دست آمده در جدول ۴ نشان می‌دهد که همه روش‌های جدید ارزیابی شده از عملکرد رضایت‌بخشی برخوردار هستند و در مقایسه با نتایج روش‌های یادگیری معمولی در جدول ۳، در هر سه معیار دقت، ویژگی و حساسیت برتری دارند. با این وجود، مدل پیشنهادی در همه معیارها نسبت به سایر روش‌های جدید مقایسه شده برتری دارد. مدل پیشنهادی با دقت ۹۰/۳۱

خون بالا و میزان مصرف نمک می‌تواند به عنوان یک ابزار مفید و زود هنگام در تشخیص فشار خون بالا در کودکان، از پیامدها و هزینه‌های ناشی از این عارضه بکاهد و گام بزرگی در مبارزه با فشار خون بالا باشد. این مدل می‌تواند به عنوان یک سیستم هشدار اولیه برای کودکان مبتلا به پیش فشار خون بالا و یا پرفشاری خون عمل کند. ترکیب روش‌های یادگیری در مدل پیشنهادی، افزایش زمان اجرایی و پیچیدگی محاسباتی را به همراه دارد، با این حال، در کاربردهای حیاتی از قبیل پزشکی که دقت و اعتماد از اهمیت و اولویت بالاتری برخوردار است، به کارگیری مدل ترکیبی مفید خواهد بود.

برای پژوهش‌های آینده، توسعه مدل پیشنهادی به همراه یک رابط کاربری مناسب برای استفاده در گوشی‌های تلفن هوشمند به طوری که والدین بتوانند با ورود اطلاعات خواسته شده، از وضعیت فشار خون کودکان خود مطلع شده و در صورت وجود خطر برای بررسی دقیق‌تر به پزشک مراجعه کنند، مورد توجه است. داده‌های مورد مطالعه در این پژوهش محدود به کودکان دبستانی شهرستان کاشمر بوده است. تعمیم مدل پیشنهادی جهت تشخیص و پیش‌بینی فشار خون در کودکان سایر مناطق و شهرها نیز به عنوان تحقیق تکمیلی پیشنهاد می‌شود. جداسازی داده‌های آزمایش و سپس متوازن‌سازی داده‌ها و بررسی نتایج به دست آمده در مدل پیشنهادی، نیز به عنوان کار آتی مورد توجه می‌باشد. همچنین بررسی ویژگی‌های دیگری از کودکان از قبیل اطلاعات بیشتری درباره سبک زندگی آن‌ها در مدل پیشنهادی و تأثیر و اهمیت هر یک از ویژگی‌ها بر ابتلای کودکان به فشارخون بالا مفید خواهد بود.

تعارض منافع

این مطالعه دارای کد اخلاق از کمیته اخلاق در پژوهش‌های زیست پزشکی دانشگاه علوم پزشکی مشهد به شماره IR.MUMS.REC.1401.038 می‌باشد. در مطالعه حاضر، نویسندگان هیچ‌گونه تعارض منافی نداشته‌اند. این مقاله حاصل تحقیق مستقل بدون حمایت مالی می‌باشد.

References

1. Chang W, Liu Y, Xiao Y, Xu X, Zhou S, Lu X, et al. Probability analysis of hypertension-related symptoms based on XGBoost and clustering algorithm. *Applied Sciences* 2019;9(6):1215. <https://doi.org/10.3390/app9061215>
2. Chowdhury MZ, Naeem I, Quan H, Leung AA, Sikdar KC, O'Beirne M, et al. Prediction of hypertension using traditional regression and machine

learning models: A systematic review and meta-analysis. *PloS one* 2022;17(4):e0266334. <https://doi.org/10.1371/journal.pone.0266334>

3. Martinez-R'ios E, Montesinos L, Alfaro-Ponce M, Pecchia L. A review of machine learning in hypertension detection and blood pressure estimation based on clinical and physiological data. *Biomedical Signal Processing and Control* 2021;68:102813. doi: 10.1016/j.bspc.2021.102813

درصد، حساسیت ۸۰/۶۵ درصد و ویژگی ۹۳/۵۴ درصد بهتر از سایر روش‌ها توانسته است تشخیص و پیش‌بینی فشار خون بالا را انجام دهد. پس از مدل پیشنهادی، روش Fang و همکاران [۴] که مبتنی بر ترکیب نتایج خروجی دو روش یادگیری KNN و LightGBM می‌باشد، کارایی بهتری را بر روی مجموعه داده‌ی متوازن شده دارد. روش Fitriyani و همکاران [۱۶] که نتایج خروجی سه روش یادگیری MLP، SVM و DT را بر اساس رویکرد رگرسیون لجستیک ترکیب می‌کند و همچنین روش Zhao و همکاران [۱۱] که مبتنی بر روش جنگل تصادفی می‌باشد، میزان دقت، ویژگی و حساسیت تقریباً نزدیک به هم دارند. در روش جنگل تصادفی Zhao و همکاران [۱۱]، نتایج خروجی چند درخت تصمیم متفاوت و مستقل بر اساس روش رأی‌گیری اکثریت با هم ترکیب می‌شوند. در بین روش‌های جدید مقایسه شده در جدول ۴، روش Chai و همکاران [۱۳] که توسعه یافته روش یادگیری MLP می‌باشد، در همه معیارها پایین‌ترین عملکرد را دارد. دلیل این عملکرد ضعیف‌تر می‌تواند به این دلیل باشد که این روش از رویکرد ترکیبی استفاده نمی‌کند و متوازن‌سازی داده‌ها نیز در آن انجام نمی‌شود. برتری دیگر مدل پیشنهادی این است که بر خلاف سایر روش‌های مقایسه شده، پیش‌بینی را به صورت جزئی‌تر انجام می‌دهد. به طوری که نمونه ورودی را به یکی از چهار دسته فشارخون طبیعی، پیش فشار خون بالا، فشارخون بالای مرحله ۱ و فشارخون بالای مرحله ۲ دسته‌بندی می‌کند. در حالی که سایر روش‌های جدید مقایسه شده، نمونه ورودی را به یکی از دو دسته طبیعی و فشار خون بالا دسته‌بندی می‌کنند. با توجه به نتایج به دست آمده، مدل ترکیبی پیشنهادی بهتر می‌تواند پیش‌بینی و تشخیص فشار خون بالا در کودکان را انجام داده و باعث بهبود دقت و کاهش میزان اشتباه گردد. این مدل با استفاده از اطلاعات اولیه دموگرافیک، داده‌های آنروپومتریک ساده و اطلاعات سبک زندگی که به راحتی از کودکان قابل یافتنی است از قبیل سن، جنسیت، قد، وزن، نبض، سابقه مصرف داروهای مؤثر بر فشار خون، سابقه بیماری‌های کلیوی، دیابت و قلبی عروقی، سابقه خانوادگی فشار

4. Fang M, Chen Y, Xue R, Wang H, Chakraborty N, Su T, et al. A hybrid machine learning approach for hypertension risk prediction. *Neural Computing and Applications* 2021;1. doi:10.1007/s00521-021-06060-0
5. Haugg F, Elgendi M, Menon C. Assessment of Blood Pressure Using Only a Smartphone and Machine Learning Techniques: A Systematic Review. *Front Cardiovasc Med* 2022; 9: 894224. doi: 10.3389/fcvm.2022.894224
6. Oanh TT, Tung NT. Predicting Hypertension Based on Machine Learning Methods: A Case Study in Northwest Vietnam. *Mobile Netw Appl* 2022; 27:2013–23. <https://doi.org/10.1007/s11036-022-01984-w>
7. Akbari M, Moosazadeh M, Ghahramani S, Tabrizi R, Kolahdooz F, Asemi Z, et al. High prevalence of hypertension among Iranian children and adolescents: a systematic review and meta-analysis. *Journal of hypertension* 2017;35(6):1155-63. doi: 10.1097/HJH.0000000000001261
8. Chai SS, Goh KL, Cheah WL, Chang YH, Ng GW. Hypertension Prediction in Adolescents Using Anthropometric Measurements: Do Machine Learning Models Perform Equally Well?. *Applied Sciences* 2022;12(3):1600. <https://doi.org/10.3390/app12031600>
9. Islam SM, Talukder A, Awal MA, Siddiqui MM, Ahamad MM, Ahammed B, et al. Machine Learning Approaches for Predicting Hypertension and Its Associated Factors Using Population-Level Data From Three South Asian Countries. *Front Cardiovasc Med* 2022; 9: 839379. doi: 10.3389/fcvm.2022.839379
10. Ghaderi Niri A, Farajzadeh N, Baybordi E. A Case Study of the Impact of Parental Diseases on the Probability of Hypertension Using Data Mining Techniques. *Journal of Health and Biomedical Informatics* 2021;7(4):354-67. [In Persian]
11. Zhao H, Zhang X, Xu Y, Gao L, Ma Z, Sun Y, et al. Predicting the risk of hypertension based on several easy-to-collect risk factors: a machine learning method. *Frontiers in Public Health* 2021;9:619429. <https://doi.org/10.3389/fpubh.2021.619429>
12. Dehghandar M, Hassani BA, Dadkhah M, Qorbani M, Kelishadi R. Diagnosis of Obesity and Hypertension in Isfahani Students Using Artificial Neural Network. *Journal of Health and Biomedical Informatics*; 2021; 8 (1):12-23.[In Persian]
13. Chai SS, Cheah WL, Goh KL, Chang YH, Sim KY, Chin KO. A Multilayer Perceptron Neural Network Model to Classify Hypertension in Adolescents Using Anthropometric Measurements: A Cross-Sectional Study in Sarawak, Malaysia. *Computational and Mathematical Methods in Medicine* 2021;2021. <https://doi.org/10.1155/2021/2794888>
14. Raza K. Improving the prediction accuracy of heart disease with ensemble learning and majority voting rule. *U-Healthcare Monitoring Systems* 2019; 179-96. <https://doi.org/10.1016/B978-0-12-815370-3.00008-6>
15. Tahmasbi H, Jalali M, Shakeri H. An Expert System for Heart Disease Diagnosis Based on Evidence Combination in Data Mining. *Journal of Health and Biomedical Informatics* 2017;3(4):251-8. [In Persian]
16. Fitriyani NL, Syafrudin M, Alfian G, Rhee J. Development of disease prediction model based on ensemble learning approach for diabetes and hypertension. *IEEE Access* 2019;7:144777-89. doi: 10.1109/ACCESS.2019.2945129
17. Rao G. Diagnosis, epidemiology, and management of hypertension in children. *Pediatrics* 2016;138(2): e20153616. <https://doi.org/10.1542/peds.2015-3616>
18. Wang YC, Cheng CH. A multiple combined method for rebalancing medical data with class imbalances. *Computers in Biology and Medicine* 2021;134:104527. <https://doi.org/10.1016/j.compbimed.2021.104527>
19. Chawla NV, Bowyer KW, Hall LO, Kegelmeyer WP. SMOTE: synthetic minority over-sampling technique. *Journal of artificial intelligence research*. 2002;16:321-57. doi: <https://doi.org/10.1613/jair.953>
20. Ijaz MF, Alfian G, Syafrudin M, Rhee J. Hybrid prediction model for type 2 diabetes and hypertension using DBSCAN-based outlier detection, synthetic minority over sampling technique (SMOTE), and random forest. *Applied Sciences* 2018;8(8):1325. <https://doi.org/10.3390/app8081325>
21. AlKaabi LA, Ahmed LS, Al Attiyah MF, Abdel-Rahman ME. Predicting hypertension using machine learning: Findings from Qatar Biobank Study. *Plos one* 2020;15(10):e0240370. <https://doi.org/10.1371/journal.pone.0240370>
22. Bouziane H, Messabih B, Chouarfia A. Profiles and majority voting-based ensemble method for protein secondary structure prediction. *Evolutionary Bioinformatics* 2011;7:EBO-S7931. <https://doi.org/10.4137/EBO.S79>
23. Chang W, Liu Y, Xiao Y, Yuan X, Xu X, Zhang S, et al. A machine-learning-based prediction method for hypertension outcomes based on medical data. *Diagnostics* 2019;9(4):178. <https://doi.org/10.3390/diagnostics9040178>
24. Lee J, Lee W, Park IS, Kim HS, Lee H, Jun CH. Risk assessment for hypertension and hypertension complications incidences using a Bayesian network. *IIE Transactions on Healthcare Systems Engineering* 2016;6(4):246-59. <https://doi.org/10.1080/19488300.2016.1232767>
25. Sokolova M, Lapalme G. A systematic analysis of performance measures for classification tasks. *Information Processing & Management* 2009;45(4):427-37. <https://doi.org/10.1016/j.ipm.2009.03.002>