



CrossMark

Original article

Design an Optimized Deep Learning Model for Liver Disease Prediction

Davoud Jalalian¹, Safoura Ashoori^{2*}

1. PhD in Electrical and Electronic Engineering, Assistant Professor, Department of Electrical Engineering, Islamic Azad University, Fouman and Shaft Branch, Fouman and Shaft, Iran

2. PhD Student in Computer Science, Rasht Branch, Islamic Azad University, Rasht, Iran

ARTICLE INFO:

Article History:

Received: 3 Aug 2024

Accepted: 6 May 2025

Published: 21 Jun 2025

*Corresponding Author:

Safoura Ashoori

Email:

safoura.ashoori@iau.ir

Citation: Jalalian D, Ashoori S. Design an Optimized Deep Learning Model for Liver Disease Prediction. Journal of Health and Biomedical Informatics 2025; 12(1): 1-17. [In Persian]

Abstract

Introduction: The incidence of liver disease resulting from excessive alcohol consumption, inhalation of polluted gases, drug use, contaminated food, and food packaging is rapidly increasing. As this disease often presents no symptoms or complications in its early stages, timely diagnosis can be challenging. However, if diagnosed early, treatment can be administered more easily and at a lower cost. In more advanced stages, liver diseases can progress to cirrhosis and liver cancer, leading to significantly more difficult treatment options and, in some cases, resulting in the patient's death.

Method: Considering the challenges involved, early diagnosis of liver disease is crucial for determining treatment solutions, duration of treatment, and recovery. With advancements in machine learning and deep learning technologies, which can analyze and learn from complex and large datasets, these tools can be employed for the early prediction of liver disease. In this study, we utilized a dataset of liver disease patients from India. After data preprocessing, a deep learning model optimized by the metaheuristic algorithm of the Lizard Search was proposed to enhance prediction accuracy by leveraging the advantages of metaheuristic algorithms and to assist in timely disease diagnosis.

Results: Our proposed method employs a deep neural network optimized using the metaheuristic RSA algorithm. This model achieved promising results in diagnosing liver disease, with an approximate accuracy of 96.9%, a reliability of around 97.2%, and an F1 score close to 96.7%. This approach demonstrates high efficiency in the early and accurate detection of liver diseases, contributing to improved treatment processes.

Conclusion: The results demonstrated that combining a feedforward neural network with metaheuristic algorithms significantly improved the accuracy and reliability of liver disease diagnosis. The proposed model enhanced prediction accuracy and diagnostic validity compared to previous methods. These advancements can effectively facilitate the early detection of liver diseases, leading to improved treatment processes and a reduction in complications associated with disease progression.

Keywords: Liver Disease, Machine Learning, Deep Learning, Metaheuristic Algorithm, Deep Neural Network

© 2025 The Author(s); Published by Kerman University of Medical Sciences. This is an open-access article distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited



CrossMark

مقاله پژوهشی

طراحی یک مدل یادگیری عمیق بهینه برای پیش‌بینی بیماری کبد

داود جلالیان^۱، صفورا عاشوری^{۲*}

۱. دکترای برق الکترونیک، استادیار، گروه مهندسی برق، دانشگاه آزاد اسلامی، واحد فومن و شفت، فومن و شفت، ایران

۲. دانشجوی دکتری، مهندسی کامپیوتر، دانشگاه آزاد اسلامی، واحد رشت، رشت، ایران

اطلاعات مقاله

سابقه مقاله

دریافت: ۱۴۰۳/۵/۱۳

پذیرش: ۱۴۰۴/۲/۱۶

انتشار برخط: ۱۴۰۴/۳/۳۱

*نویسنده مسئول:

صفورا عاشوری

ایمیل:

safoura.ashoori@iau.ir

ارجاع: جلالیان داود، عاشوری

صفورا. طراحی یک مدل یادگیری عمیق بهینه برای پیش‌بینی بیماری کبد. مجله انفورماتیک سلامت و زیست پزشکی ۱۴۰۴؛ ۱۲(۱): ۱-۱۷.

مقدمه: ابتلا به بیماری کبد به دلیل نوشیدن بیش از حد الکل، استنشاق گازهای آلوده، داروها، مواد غذایی آلوده و بسته‌بندی غذا به سرعت در حال افزایش است. با توجه به این که این بیماری در مراحل اولیه عوارض و علائمی ندارد تشخیص به موقع آن کار دشواری می‌باشد. اگر این بیماری در مراحل اولیه تشخیص داده شود درمان آن به آسانی و با صرف هزینه کمتر امکان‌پذیر می‌باشد. بیماری‌های کبدی در مراحل پیشرفته‌تر تبدیل به سیروز و سرطان کبدی می‌گردد که به مراتب راه‌های درمانی بسیار دشوارتری را به دنبال دارند حتی در بعضی موارد منجر به مرگ بیمار می‌گردد.

روش کار: با توجه به چالش‌های مطرح شده تشخیص زود هنگام این بیماری در ارائه راهکارهای درمانی، طول درمان و بهبودی بسیار مهم است. با پیشرفت فناوری یادگیری ماشینی و یادگیری عمیق و توانایی آن در تحلیل و یادگیری ویژگی از مجموعه داده‌های پیچیده و حجیم می‌توان از این ابزار جهت پیش‌بینی زود هنگام بیماری کبد استفاده نمود، بدین منظور در پژوهش حاضر، از مجموعه داده بیماران کبدی هند استفاده شد. پس از پیش پردازش داده‌ها، یک مدل یادگیری عمیق که به وسیله الگوریتم فرایبتکاری جستجوی خزندگان بهینه شده مطرح گردید تا با استفاده از مزیت الگوریتم‌های فرایبتکاری دقت پیش‌بینی را افزایش دهد و به تشخیص به موقع بیماری کمک کند.

یافته‌ها: در مطالعه حاضر، روش پیشنهادی شامل شبکه عصبی عمیق بهینه‌شده با بهره‌گیری از الگوریتم فرایبتکاری RSA است. این مدل توانست در تشخیص بیماری کبدی، با دقت تقریبی ۹۶/۹ درصد، صحت حدود ۹۷/۲ درصد و امتیاز F1 نزدیک به ۹۶/۷ درصد، نتایج قابل قبولی را نشان دهد. این رویکرد نشان‌دهنده کارایی بالا در تشخیص زود هنگام و دقیق بیماری‌های کبدی است که می‌تواند به بهبود فرآیندهای درمانی کمک کند.

نتیجه‌گیری: نتایج نشان داد که استفاده از شبکه عصبی پیشخور و الگوریتم‌های فرایبتکاری، موجب بهبود قابل توجهی در دقت و صحت تشخیص بیماری کبدی شد. مدل پیشنهادی توانست به‌طور قابل ملاحظه‌ای دقت پیش‌بینی و صحت نتیجه‌گیری را نسبت به روش‌های پیشین افزایش دهد. این پیشرفت‌ها می‌تواند در تشخیص زود هنگام و به موقع بیماری‌های کبدی نقش مؤثری ایفا کند و به بهبود روند درمان و کاهش عوارض ناشی از پیشرفت بیماری کمک نماید.

کلیدواژه‌ها: بیماری کبد، یادگیری ماشین، یادگیری عمیق، الگوریتم فرایبتکاری، شبکه عصبی عمیق

مقدمه

امروزه بخش بهداشت و درمان نقش مهمی در اقتصاد کشورها ایفا می‌کند و با اقدامات ضروری به بهبود سلامت جامعه کمک می‌کند. این بخش شامل مراقبت‌های اولیه، ثانویه و سوم است و هدف اصلی آن ارائه خدمات با کیفیت بالا و پیشگیری از بیماری‌ها در مراحل اولیه است [۱]. تشخیص زودهنگام بیماری‌های کبدی برای جلوگیری از نارسایی کامل کبد، که می‌تواند منجر به مرگ شود، بسیار حیاتی است [۲،۳]. کبد، بزرگ‌ترین عضو بدن انسان، نقش حیاتی در سلامت کلی بدن دارد. این عضو مسئول عملکردهای متابولیکی کلیدی از جمله تبدیل مواد مغذی رژیم غذایی به مواد قابل استفاده، ذخیره و تأمین این مواد برای سلول‌ها، تولید صفرا، پردازش هموگلوبین، پاکسازی خون، تبدیل مواد سمی به مواد بی‌ضرر و تقویت ایمنی بدن است [۴،۵]. با این حال، به دلیل سبک زندگی ناسالم، بسیاری از افراد در سراسر جهان از بیماری‌های کبدی، از مشکلات حاد تا شدید، رنج می‌برند [۶]. بیماری‌های کبدی شامل سیروز، هپاتیت ویروسی، کبد چرب، بیماری‌های ژنتیکی کبد و بیماری کبدی الکلی هستند. سیروز کبدی عامل حدود ۱/۸ درصد مرگ‌ومیرها در اروپا است و در دهه اخیر به دلیل افزایش مصرف الکل، عفونت‌های مزمن هپاتیت و بیماری‌های کبدی مرتبط با چاقی، شیوع آن به‌طور قابل توجهی افزایش یافته است. تحقیقات اخیر تفاوت‌های جنسیتی در شیوع، تشخیص و مدیریت بیماری‌های کبدی را نشان داده‌اند پیشرفت‌های اخیر در استفاده از هوش مصنوعی در مراقبت‌های بهداشتی، امکان بهره‌گیری از مجموعه داده‌های بالینی بزرگ را برای یادگیری ماشین فراهم کرده است. طبقه‌بندی یادگیری ماشین در مقایسه با مدل‌های آماری سنتی عملکرد بهتری دارند [۷]. در این پژوهش، برای تشخیص زودهنگام نارسایی کبد از یک مدل یادگیری عمیق مبتنی بر شبکه عصبی عمیق استفاده خواهد شد.

Jin و همکاران [۸]، مفهوم تکنیک‌های طبقه‌بندی مختلفی را توضیح دادند که به پزشکان کمک می‌کند تا بیماری را سریع و کارآمد تشخیص دهند. طبقه‌بندی‌های مختلفی مانند بیزین ساده (Naïve Bayes)، پرسپترون چند لایه (Multilayer Perceptron)، درخت تصمیم (Decision Tree) و (K(KNN K-Nearest Neighbors) نزدیک‌ترین همسایه بر اساس چندین پارامتر مانند ویژگی (Feature)، حساسیت (Recall) و غیره مقایسه و تجزیه و تحلیل شدند. الگوریتم‌ها با استفاده از ابزار این سورس وکا (Weka) پیاده‌سازی شدند و مجموعه داده‌ها از مخزن UCI (University of California, Irvine) جمع‌آوری شد. نتایج تجربی نشان داد که از نظر دقت، بیزین ساده نتایج طبقه‌بندی بهتری را دارد، در حالی که رگرسیون لجستیک و جنگل تصادفی نتایج بهتری از نظر دو معیار یادآوری و حساسیت به‌دست آوردند. Singh و همکاران در مطالعه‌ای [۹] از الگوریتم‌های طبقه‌بندی مختلف مانند رگرسیون لجستیک، ماشین بردار پشتیبان و نزدیک‌ترین همسایه برای پیش‌بینی بیماری کبد استفاده کردند. مقایسه همه این الگوریتم‌ها بر اساس درصد صحت طبقه‌بندی که از طریق ماتریس پیچیدگی پیدا می‌شود، انجام شد. در این کار، رگرسیون لجستیک و K-نزدیک‌ترین همسایه به بالاترین درصد صحت دست یافته‌اند؛ اما رگرسیون لجستیک بالاترین حساسیت را داشت. درصد صحت به‌دست آمده برای ماشین بردار پشتیبان ۷۱/۹۷٪ و برای رگرسیون لجستیک و K-نزدیک‌ترین همسایه ۷۳/۹۷٪ بود؛ بنابراین می‌توان نتیجه گرفت که رگرسیون لجستیک برای پیش‌بینی بیماری کبد مناسب است. Pathan و همکاران در مطالعه‌ای [۱۰] الگوریتم‌های طبقه‌بندی مختلفی مانند بیزین ساده، J48، کیسه‌کشی (Bagging)، جنگل تصادفی (Random Forest) و آدا بوست (Ada Boost) استفاده کردند. همچنین از روش پیش پردازش برای تقسیم داده‌ها به دو گروه بیماران کبدی و بیماران غیر کبدی استفاده شد که با استفاده از الگوریتم خوشه‌بندی K means انجام شد. پیاده‌سازی الگوریتم‌های طبقه‌بندی مختلف با استفاده از ابزار این سورس وکا انجام شد. این الگوریتم‌ها بر اساس معیارهای ارزیابی مانند صحت، میزان خطا و غیره با هم مقایسه شدند. پس از انجام مقایسه، مطالعه تطبیقی نشان داد که جنگل تصادفی نتایج بهتری را در مقایسه با سایر الگوریتم‌ها ارائه می‌دهد. Adil و همکاران در مطالعه‌ای [۱۱]، از رویکرد یادگیری ماشین برای طبقه‌بندی باینری بیمار کبدی (یعنی بیمار کبدی یا غیر کبدی) با استفاده از جنسیت بیمار و داده‌های نمونه پزشکی آزمایشگاهی استفاده کردند. مجموعه داده مورد استفاده در مخزن یادگیری ماشینی UCI با عنوان «سوابق بیماران کبدی هند» بود. هدف این مطالعه استفاده از تکنیک طبقه‌بندی محاسباتی ساده مانند رگرسیون لجستیک و مقایسه نتایج آن با نتایج قبلی به‌دست‌آمده بر روی مجموعه داده مورد نظر، توسط محققان دیگر است. نتایج طبقه‌بندی رگرسیون لجستیک اهمیت خود را در این مجموعه داده با دستیابی به صحت بالاتر ۷۴٪ در طبقه‌بندی نسبت به طبقه‌بند ساده بیز، C5.0، درخت تصمیم (Decision Tree)،



ماشین بردار پشتیبان، شبکه عصبی مصنوعی و k نزدیک ترین همسایه ثابت کرد. Baitharu و Pani [۱۲]، رویکردی برای تشخیص اختلالات کبدی از طریق تجزیه و تحلیل مجموعه داده‌های اختلالات کبدی ارائه کردند. چندین الگوریتم مانند بیزین ساده، شبکه عصبی چند لایه، ZeroR (Zero Rule algorithm)، JBK (Instance-Based Learner with parameter)، VFI J48 (Feature Importance Versione ejra shodeye algorithm C4.5 dar WEKA) و پرسپترون چندلایه در این پژوهش مورد استفاده قرار گرفتند. مجموعه داده‌های مورد استفاده از پایگاه داده UCI جمع‌آوری شد. الگوریتم‌ها با استفاده از ابزار وکا پیاده سازی شدند. نتایج تجربی نشان داد که پرسپترون چندلایه نتایج طبقه‌بندی بهتری را در مقایسه با سایر الگوریتم‌ها ارائه می‌دهد. بنابراین، پرسپترون چندلایه می‌تواند بیشتر برای تشخیص اختلال کبدی به طور مؤثر مورد استفاده قرار گیرد. Geetha و همکاران در پژوهش خود [۱۳] با استفاده از تکنیک‌های یادگیری ماشینی (ML (Machin Learning)، روش‌های تشخیص بیماری کبدی در بیماران پیشنهاد و ارزیابی کردند. ماشین بردار پشتیبان، رگرسیون لجستیک، شامل دو تکنیک اصلی یادگیری ماشین هستند که در این پژوهش استفاده شدند. برای ارزیابی الگوریتم‌ها پایگاه داده UCI به کار گرفته شد. با استفاده از تمامی مدل‌ها، تحلیل پیش‌بینی اجرا شده و عملکرد آن‌ها ارزیابی شده است. برای طبقه‌بند ماشین بردار پشتیبان ۷۵/۰۴ درصد صحت و برای رگرسیون لجستیک ۷۳/۲۳ درصد صحت گزارش شده است. Vijayarani و Dhayanand [۱۴] یک رویکرد پیش‌بینی اختلال کبدی را با استفاده از الگوریتم‌های طبقه‌بندی مختلف ارائه دادند. در این رویکرد از الگوریتم‌های طبقه‌بندی بیزین ساده و ماشین بردار پشتیبان استفاده کردند. این دو الگوریتم بر اساس پارامترهای عملکردی که شامل معیارهای صحت طبقه‌بندی و معیارهای زمان اجرا می‌شود، مقایسه شدند. سیستم پیشنهادی با استفاده از ابزار MATLAB نسخه ۲۰۱۳ پیاده‌سازی شد و بر مجموعه داده‌های جمع‌آوری شده از پایگاه داده UCI اعمال گردید. در نتایج تجربی، مشاهده شد که ماشین بردار پشتیبان به دلیل بالاترین صحت طبقه‌بندی، از الگوریتم بیزین ساده بهتر عمل کرد و می‌توان از آن بیشتر در پیش‌بینی بیماری‌های کبدی استفاده کرد. Abdar و همکاران جهت تشخیص زودهنگام و کارآمد نارسایی کبد از یک رویکرد مدل‌سازی یادگیری ماشین یکپارچه استفاده کردند [۱۵]؛ این رویکرد شامل ادغام یک شبکه عصبی چندلایه پرسپترون با درخت‌های تصمیم‌گیری مختلف مانند (C5.0) [۱۲]، CHAID [۱۱]، طبقه‌بند و درخت رگرسیون تقویت‌شده [۱۳] بود. در این پژوهش از مجموعه داده‌های مخزن دانشگاه کالیفرنیا، ایروین (UCI) استفاده کردند. چندین معیار ارزیابی مانند حساسیت، ویژگی، دقت و غیره مورد بررسی قرار گرفتند. نتایج نشان می‌دهد که تلفیق شبکه عصبی (Multilayer MLP (Perceptron و روش‌های C5.0 با دستیابی به حساسیت ۹۴/۱۶٪ و ویژگی ۹۳/۷۵٪ در مقایسه با سایر الگوریتم‌ها از عملکرد بهتری برخوردار بود.

مطالعه اخیر اینتل بر روی پیاده‌سازی طبقه‌بندی مجموعه داده بیماران کبدی هند با استفاده از توزیع پایتون برای Intel® بر روی پردازنده Intel® Xeon® Scalable تمرکز دارد و از روش‌های مختلف پیش‌پردازش استفاده شد و تأثیر این روش‌ها در طبقه‌بندی مورد بررسی قرار گرفت. همچنین به کمک ویژگی‌های مختلف، طبقه‌بندی بیماران کبدی با هدف پیش‌بینی این که آیا یک فرد به بیماری کبدی مبتلا است یا خیر انجام شد و با استفاده از SMOTE (Synthetic Minority Oversampling Technique) به عنوان روش پیش‌پردازش، روی داده‌های کبدی UCI، که ممکن است مانند بسیاری از مجموعه داده‌های دنیای واقعی، با مشکل عدم تعادل کلاس مواجه شوند، نتایج خوبی به دست آمد [۱۶]. به این معنی که تعداد نمونه‌های مربوط به یک نوع خاص از حمله یا دفاع (یا هر ویژگی دیگری که به عنوان کلاس در نظر گرفته شود) ممکن است به طور قابل توجهی کمتر از سایر انواع باشد، SMOTE با ایجاد نمونه‌های مصنوعی جدید برای کلاس اقلیت، به تعادل داده‌ها کمک می‌کند. این کار با انتخاب نمونه‌های موجود از کلاس اقلیت و ایجاد نمونه‌های جدید در امتداد خطوطی که این نمونه‌ها را به هم وصل می‌کنند، انجام می‌شود. این عمل کمبود داده را جبران نموده و موجب تعادل بین کلاس‌ها شده و در بهبود دقت تأثیرگذار خواهد بود. همچنین از الگوریتم جنگل تصادفی به عنوان طبقه‌بند استفاده شد. در این کار صحت ۷۴/۶٪ برای جنگل تصادفی و ۷۱/۲٪ برای ماشین بردار پشتیبان به دست آمد [۱۷]. در پژوهش منتظری ۵ مدل یادگیری

¹ Intel Corporation, *Intel® Xeon® Scalable Processor 8280*, 2019. <https://www.intel.com/content/www/us/en/products/sku/120496/intel-xeon-platinum-8280-processor-38-5m-cache-2-70-ghz/specifications.html>

ماشین بیز ناو (Bayes Naïve)، بردار ماشین پشتیبان (SVM(Support Vector Machine)، آدبوست (ADBOOST) جنگل تصادفی (Trees Random Forest) و INN به منظور تشخیص بیماری کبد استفاده شده است. در مدل دسته‌بند جنگل تصادفی تعداد درختان به کار برده شده ۱۰ درخت است و در مدل آدبوست تعداد مراحل اجرا ۱۰ مرحله است. ۵ مدل ماشین یادگیری اجرا شده‌اند، این ۵ مدل از نظر ویژگی، حساسیت، سطح زیر منحنی راک و دقت دسته‌بندی مقایسه شدند، از نظر دقت دسته‌بندی مدل یادگیری ماشین جنگل تصادفی، با دقت ۷۲/۲۱٪ دارای دقت بهینه است. دو مدل جنگل تصادفی، شبکه بیزین (Naïve Bayes) نسبت به سایر مدل‌های یادگیری ماشین در موقعیت بهتری قرار دارند و در نتیجه سطح زیر منحنی راک بیشتری دارند [۱۸]. در پژوهش خامنه و همکاران، شناسایی بیماران مبتلا به سیروز در معرض ابتلا به سرطان کبد به کمک رویکردهای داده کاوی مطرح شد. میزان تأثیرگذاری مشخصه‌های مورد بررسی بیماران در پیش‌بینی صحیح سرطان در بیماران سیروتیک، از مزایای تحقیق آنان بود که درخت تصمیم توانست میزان دقت ۸۸ درصدی را برای بیماران با اتیولوژی ویروسی و دقت ۹۲ درصدی را برای بیماران با پاتیولوژی غیر ویروسی ارائه نماید [۱۹]. تشخیص زود هنگام و دقیق بیماری‌های کبدی به دلیل تأثیر قابل توجه آن در بهبود روند درمان، کاهش عوارض و استفاده از روش‌های درمانی ساده‌تر و کم‌هزینه‌تر، از اهمیت ویژه‌ای برخوردار است. این پژوهش با بهره‌گیری از روش‌های پیشرفته یادگیری عمیق و یادگیری ماشین، به منظور شناسایی بیماری کبدی در مراحل اولیه با دقت بالا انجام شده است تا گامی مؤثر در جهت بهبود فرآیند تشخیص و مدیریت این بیماری برداشته شود.

روش کار

در این مطالعه برای تشخیص نارسایی کبد از یک مدل یادگیری عمیق بر پایه شبکه عصبی عمیق استفاده شد. بدین منظور از مجموعه داده (Indian Liver Pati) ILPD موجود در پایگاه داده یادگیری ماشین UCI استفاده شد. این مجموعه داده، از شمال شرقی آندرا پرداش (Andhra Pradesh)، هند جمع‌آوری شده و شامل مجموعه‌ای از علائم بالینی است. بعد از اعمال روش‌های مربوط به پیش پردازش بر روی مجموعه داده موجود که هدف از پیش پردازش، حذف داده‌های پرت و پاک‌سازی داده‌ها (Clean data) و تجزیه و تحلیل اطلاعات بیماران از مجموعه داده و استخراج ویژگی‌ها و داده‌های مؤثر در تشخیص بیماری با استفاده از روش محاسبه میانگین، میان، انحراف معیار، ماکزیمم مینیمم و واریانس برای هر ویژگی عددی می‌باشد، این ویژگی‌های پایه، اطلاعات اولیه‌ای در مورد توزیع هر ویژگی ارائه می‌دهند؛ به عنوان مثال میانگین سطح آلبومین سرم می‌تواند نشان دهنده عملکرد کلی کبد باشد. سپس از کل مجموعه داده، داده‌ها به صورت تصادفی به دو مجموعه آموزش و آزمون تقسیم می‌شوند. این به این معنا است که داده‌ها به صورت تصادفی انتخاب شدند تا از هرگونه تعصب جلوگیری شود و نمونه‌گیری تصادفی به بهترین شکل نمایانگر کلیت داده‌ها باشد، ۳۰ درصد داده‌ها برای آزمون در نظر گرفته شدند و بقیه یعنی ۷۰ درصد برای آموزش استفاده شدند. به منظور جلوگیری از تکرار داده و بیش برآزش (overfitting) و اطمینان از این که داده‌های آزمون و آموزش مستقل از هم هستند، استفاده از روش‌های اعتبارسنجی متقابل k تایی (k -fold cross-validation) استفاده شد. با توجه به نامتعادل بودن داده‌ها برای جلوگیری از بیش برآزش و تمایل به کلاس خاصی برای پیش‌بینی از روش ارزیابی متقاطع K که در آن نمونه اصلی به صورت تصادفی به ۱۰ زیرنمونه تقسیم می‌شود، استفاده شده که در آن K برابر با ۱۰ می‌باشد که در آن داده‌ها به ۱۰ زیرمجموعه تقسیم می‌شوند برای هر بخش یکی از زیرمجموعه‌ها برای آزمون و سایر زیرمجموعه‌ها برای آموزش استفاده می‌شوند مدل بر روی داده‌های آموزشی، آموزش داده می‌شود و بر روی داده‌های آزمون ارزیابی می‌شود و دقت آن ثبت می‌شود این کار برای هر K زیرمجموعه تکرار می‌شود به این ترتیب هر کدام از زیرمجموعه‌ها یک‌بار به عنوان داده‌های آزمایشی و باریگر به عنوان داده آموزشی استفاده می‌شود و در هر بار دقت‌هایی که از ۱۰ زیرنمونه ۹ تایی آن‌ها برای داده‌های آموزشی و یک زیرنمونه به منظور آزمایش مدل استفاده می‌شود. همچنین فرآیند ارزیابی متقاطع، ۱۰ بار تکرار می‌شود. میانگین نتایج حاصل از ۱۰ بار اجرا به عنوان نتایج نهایی مورد استفاده قرار می‌گیرد. X ماتریسی $n \times m$ و y ماتریسی $n \times 1$ می‌باشد که m معرف تعداد ویژگی‌ها می‌باشد و n تعداد دسته‌ها (بیمار، سالم) در هر مجموعه داده را نشان می‌دهد. ارزیابی متقاطع موجب می‌شود که دقت مدل بالا برود. همچنین این فرآیند اطمینان می‌دهد که همه داده‌ها به طور کامل در عملیات آموزش و آزمون قرار می‌گیرند، و می‌تواند منجر به کاهش بیش برآزش

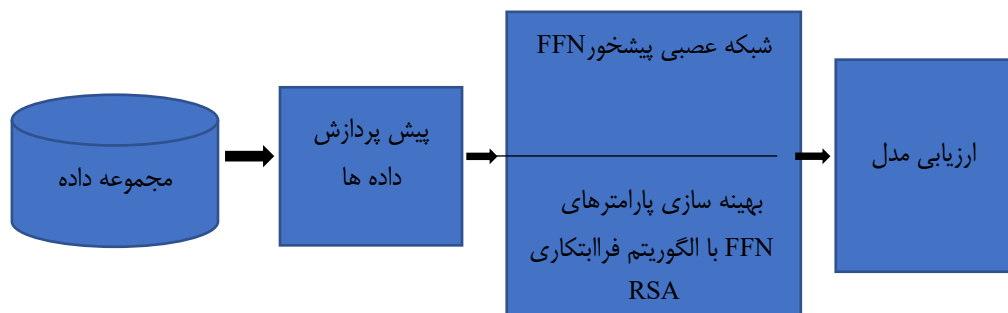


(overfitting) شود. علاوه بر این، به منظور مدیریت داده‌های نامتوازن از روش Class Weighting استفاده شد. در این روش وزن بیشتری به کلاس اقلیت (افراد غیر کبدی) اختصاص داده شد تا مدل در فرآیند یادگیری به این کلاس توجه بیشتری داشته باشد. برای این منظور، وزن کلاس‌ها به صورت معکوس نسبت به فراوانی آن‌ها تنظیم شد: وزن کلاس اقلیت (افراد غیر کبدی) برابر با $\frac{n_{\text{total}}}{n_{\text{non-liver}}}$ و وزن کلاس اکثریت (بیماران کبدی) برابر با $\frac{n_{\text{total}}}{n_{\text{liver}}}$ محاسبه شد، که در آن n_{total} تعداد کل نمونه‌ها (۵۸۳) و n_{liver} و $n_{\text{non-liver}}$ به ترتیب تعداد نمونه‌های بیماران کبدی (۴۱۶) و افراد غیر کبدی (۱۶۷) هستند. این وزن‌ها در تابع هزینه مدل اعمال شدند تا تأثیر کلاس اقلیت در فرآیند بهینه‌سازی افزایش یابد. آزمایش با همان تنظیمات 10-fold cross-validation و ۱۰ بار تکرار انجام شد تا نتایج قابل مقایسه با مدل اصلی باشند. نتایج آزمایش تکمیلی با روش class weighting نشان داد که این روش به بهبود عملکرد مدل در شناسایی کلاس اقلیت (افراد غیر کبدی) کمک کرد، در حالی که دقت کلی مدل همچنان در سطح بالایی باقی ماند. مقایسه نتایج این آزمایش با مدل class weighting با روش k-fold cross-validation در جدول ۴ آورده خواهد شد. در مرحله پیش‌پردازش داده‌ها، اقداماتی برای بهبود کیفیت داده‌ها انجام شد. مقادیر از دست رفته با استفاده از میانگین گیری حسابی جایگزین شدند و ستون جنسیت با مقادیر غیر عددی (مرد و زن) به مقادیر عددی ۱ برای مرد و ۰ برای زن تبدیل شد. برای نرمال سازی داده‌ها و یکسان سازی مقیاس ویژگی‌ها، از روش StandardScaler استفاده شد که داده‌ها را به گونه‌ای تبدیل کرد که میانگین آن‌ها صفر و واریانس آن‌ها یک شود (فرمول‌های ۱ تا ۳). این فرآیند به بهبود عملکرد مدل‌های یادگیری ماشین، به ویژه در داده‌هایی با توزیع تقریباً نرمال و مقیاس‌های متفاوت، کمک کرد. متغیرهای موجود در مجموعه داده شامل سن بیمار، جنسیت بیمار، کل بیلی روبین، بیلی روبین مستقیم، آلکالین فسفاتاز، آلانین آمینوترانسفراز، آسپارتات آمینوترانسفراز، کل پروتئین‌ها، آلبومین، نسبت آلبومین و گلوبولین می‌باشد که در جدول ۱ آورده شده است.

جدول ۱: مجموعه داده استفاده شده در مدل

ردیف	ویژگی	توضیحات
۱	Age	سن بیمار
۲	Gende	جنسیت بیمار
۳	TB (Total bilirubin)	کل بیلی روبین
۴	DB(Directbilirubin)	بیلی روبین مستقیم
۵	Alkaline Phosphatase	آلکالین فسفاتاز
۶	ALT(AlanineAminotransferas)	آلانین آمینوترانسفراز
۷	AST(AspartateAminotransferase)	آسپارتات آمینوترانسفراز
۸	TP (Total Proteins)	کل پروتئین‌ها
۹	ALB(Albumin)	آلبومین
۱۰	A/G Ratio (Albumin/Globulin Ratio)	نسبت آلبومین و گلوبولین
۱۱	Labele	فیلد انتخاب‌گر مورد استفاده برای تقسیم داده‌ها به دو مجموعه (برچسب شده توسط کارشناسان)

پیش داده‌ها وارد شبکه عصبی پیش‌خور (FFN (Feedforward Neural Network شدند. به منظور بالا بردن کارایی مدل، وزن‌های شبکه عصبی با استفاده از الگوریتم فرا ابتکاری جستجوی خزندگان (RSA (Reptile Search Algorithm بهینه شد. عملکرد این تکنیک با معیارهای مختلف مانند صحت، دقت، یادآوری ارزیابی شد. برای ارزیابی جامع عملکرد مدل پیشنهادی (FFN+RSA) در طبقه‌بندی بیماران کبدی و غیر کبدی، از منحنی ROC (Receiver Operating Characteristic) و مقدار (Area Under Curve) AUC استفاده شد. منحنی ROC نرخ مثبت صحیح (TPR) یا Recall را در مقابل نرخ مثبت کاذب (FPR) در آستانه‌های مختلف تصمیم‌گیری نشان می‌دهد. AUC معیاری از توانایی مدل در تمایز بین کلاس‌های مثبت (بیمار) و منفی (سالم) است، که مقادیر نزدیک به ۱ نشان‌دهنده عملکرد قوی مدل هستند. این معیارها با استفاده از خروجی‌های احتمالی لایه خروجی (Sigmoid) و داده‌های آزمون ۳۰٪ از مجموعه داده ILPD محاسبه شدند تا از بیش‌برازش و نشت داده جلوگیری شود. همچنین نتایج مدل با کارهای پیشین مقایسه شد. بلاک دیاگرام مدل پیشنهادی در شکل ۱ آورده شده است.



شکل ۱: بلاک دیاگرام مدل پیشنهادی

طراحی و معماری مدل پیشنهادی

یکی از اصلی‌ترین مراحل برای جستجو و کار با دیتا مرحله پیش پردازش داده می‌باشد بعد از پیش پردازش داده‌ها تبدیل به دانش خواهند شد. دانشی که برای حل مسائل و یافتن راه‌حلی برای آن‌ها مفید می‌باشد. با پیش پردازش داده‌ها و مشخص شدن داده‌های پرت (Outliers)، داده‌های از دست رفته (Missing Data)، داده‌های نویزها (Noisy Data) و... می‌توان الگوی معناداری از داده‌ها را کشف کرد و از این الگوها در راستای حل مسائل بهره برد.

مراحل پیش پردازش داده‌ها

۱- جایگزینی مقادیر از دست رفته

با بررسی در مجموعه داده مشاهده می‌شود که برخی از المان‌های سنجش بیماری دارای مقادیر از دست رفته هستند برای پر کردن فیلدهای خالی یا همان مقادیر از دست رفته از روش میانگین‌گیری حسابی (Arithmetic mean) استفاده شد و مقادیر به دست آمده در جدول جایگزین شدند.

۲- جایگزینی مقادیر غیر عدد با مقادیر عدد

در میان مجموعه داده موجود فقط ستون مربوط به جنسیت (gender) می‌باشد که دارای مقادیر غیر عددی می‌باشد برای رسیدن به نتایج بهتر مقادیر ستون جنسیت نیز عددی شد. برای عددی کردن مقادیر این ستون به‌جای مقدار Female عدد صفر و به‌جای مقدار Male عدد ۱ قرار گرفت.

برای نرمال‌سازی داده‌ها از standard scaler استفاده شد. یکی از روش‌های متداول در پیش پردازش داده‌ها است زمانی که داده‌ها دارای توزیع تقریباً نرمالی باشند و ویژگی‌های مختلف داده‌ها دارای مقیاس‌های بسیار متفاوتی باشند می‌تواند بسیار مفید باشد با توجه به ماهیت داده‌ها و مقیاس متفاوت و توزیع نرمال آن‌ها در مدل پیشنهادی از این روش استفاده شد. این روش برای مقیاس داده‌ها به گونه‌ای استفاده می‌شود که میانگین آن‌ها صفر و واریانس آن‌ها یک شود با استفاده از این روش داده‌ها به گونه‌ای تبدیل می‌شوند که مقادیر آن‌ها در یک بازه استاندارد $[-3, 3]$ قرار بگیرد این کمک می‌کند که مدل‌های مختلف یادگیری ماشین (Machine learning) بهتر و پایدارتر آموزش داده شوند.

$$Z = X - \frac{Y}{\sigma} \quad (1)$$

$$\sigma^2 = \frac{\sum_{i=1}^n (x_i - y)^2}{n - 1} \quad (2)$$

$$\sigma = \sqrt{\sigma^2} \quad (3)$$

Z نشان دهنده مقدار نرمال شده، X نشان دهنده مقدار اصلی، y مقدار میانگین داده‌ها و σ^2 معرف انحراف معیار است که از فرمول ۲ محاسبه می‌شود و σ نشان دهنده واریانس است و طبق فرمول ۳ از جذر واریانس به دست می‌آید [۲۰].

شبکه عصبی پیشخور FFN



یکی از مهم‌ترین و جدیدترین مباحث در یادگیری ماشین و یادگیری عمیق، مبحث شبکه عصبی عمیق است شبکه عصبی عمیق یک نوع از شبکه‌های عصبی مصنوعی است که شامل چندین لایه از نورون‌ها (Neuron) است. شبکه عصبی عمیق شامل لایه‌های ورودی و خروجی به همراه چندین لایه مخفی در بین آن‌ها می‌باشد این شبکه‌ها به طور خاص برای یادگیری وزن‌ها و الگوها و بررسی داده‌های پیچیده و ساختارهای عمیق استفاده می‌شود. شبکه عصبی پیشخور نوعی شبکه عصبی است که به آن پرسپترون چندلایه هم گفته می‌شود. درواقع چندین لایه پرسپترون که به هم متصل هستند یک شبکه عصبی پیشخور را تشکیل می‌دهند. یک شبکه عصبی پیشخور لایه ورودی، لایه‌های میانی یا نهان و لایه خروجی دارد که به هم متصل هستند. این لایه‌ها از چندین پرسپترون یا نورون تشکیل شده است. پرسپترون کوچک‌ترین واحد یک شبکه عصبی است. به عبارت دیگر بلاک اصلی ساخت یک شبکه عصبی را پرسپترون تشکیل می‌دهد. درواقع پرسپترون یک شبکه عصبی تک‌لایه (Single Layer Neural Network) است که به عنوان یک طبقه‌بند باینری کار می‌کند. یک پرسپترون شبکه عصبی است که فقط یک نورون دارد. اطلاعات در طول این شبکه از ورودی به سمت خروجی حرکت می‌کند و از لایه‌های مختلفی عبور می‌کند تا در نهایت خروجی مدنظر را در اختیارمان قرار دهد. در این شبکه حلقه‌ای وجود ندارد که خروجی شبکه دوباره به عنوان ورودی به خود شبکه بازگردد به طور کلی یک شبکه عصبی پیشخور از یک لایه ورودی، یک یا چند لایه نهان و یک لایه خروجی تشکیل شده است. هر یک از این لایه‌ها چندین نورون دارند. این نورون‌ها به نورون‌های لایه بعدی متصل هستند و خروجی هر نورون به نورون‌های لایه بعد وارد می‌شود تا در نهایت نورون موجود در لایه خروجی پیش‌بینی نهایی را به ما اعلام کند. برای حل یک مسئله کلاسه‌بندی باینری از شبکه عصبی پیشخور استفاده می‌شود [۲۱]: در مدل پیشنهادی با توجه به مجموعه داده‌های موجود و این که هدف، تشخیص افراد بیمار کبدی و افراد سالم می‌باشد از این شبکه عصبی عمیق استفاده شد و نتایج مطلوبی جهت تشخیص به دست آمد.

الگوریتم جستجوی خزندگان RSA

الگوریتم‌هایی فراابتکاری الگوریتم‌هایی هستند که برای حل مسائل بهینه‌سازی پیچیده و چالش‌برانگیز استفاده می‌شوند و بر اساس مفاهیم و الهام‌بخشی‌هایی از حوزه‌های مختلف مانند طبیعت، فرایندهای اجتماعی و مفاهیم ریاضی و آماری طراحی شده‌اند. الگوریتم‌های فرا ابتکاری برای حل مسائل بهینه‌سازی پیچیده معمولاً به صورت تکاملی عمل می‌کنند و بهبود تدریجی را در فضای جستجوی مسئله ایجاد می‌کنند ازجمله مزایای مهم الگوریتم‌های فرا ابتکاری می‌توان به انعطاف‌پذیری در مسائل پیچیده، قابلیت کاربرد در فضاهای جستجوی بزرگ، قابلیت مقاومت در برابر شرایط نامطلوب محیطی و قابلیت تطبیق با مسائل جدید اشاره کرد. هدف این مطالعه ارائه چارچوبی جهت پیش‌بینی بیماری کبدی با دقت بالا براساس مزیت‌های ترکیب روش‌های یادگیری عمیق و الگوریتم‌های فراابتکاری است. از این رو از الگوریتم فراابتکاری جستجوی خزندگان استفاده شد. این الگوریتم جزء الگوریتم‌های فراابتکاری بوده که در سال ۲۰۲۲ مطرح شده است [۲۲،۲۳]: RSA به دلیل ساختار اجتماعی و مکانیزم جستجوی خود می‌تواند در فضایی با قابلیت برداشت جهانی فوق‌العاده عمل کند و به راحتی از نقاط بهینه محلی فرار کند. همچنین استفاده از استراتژی‌های خود این الگوریتم می‌تواند در محیط‌های پیچیده و پر نویز به خوبی عمل کند. به این ترتیب قابلیت تعمیم مدل‌های یادگیری ماشین و بهینه‌سازی وزن‌ها را افزایش دهد. الگوریتم جستجوی خزندگان می‌تواند وزن‌های شبکه را به صورت همزمان بهینه‌سازی کند و با استفاده از تعاملات میان خزندگان، به جستجوی دقیق‌تری دست یابد. حسن استفاده از الگوریتم‌های فراابتکاری در این است که می‌توانند با بهینه کردن شبکه‌های یادگیری عمیق باعث بهبود در معیارهای پیش‌بینی شوند. با توجه به این نکات و جدید بودن الگوریتم در حوزه بهینه‌سازی، الگوریتم جستجوی خزندگان برای این تحقق مورد توجه قرار گرفت. استراتژی‌های RSA در فازهای مختلف به صورت کامل شرح داده شده است.

این الگوریتم از نحوه رفتار محاصره و شکار کروکودیل الهام گرفته شده است. الگوریتم جستجوی خزندگان با ایجاد مجموعه‌ای تصادفی از گزینه‌های حل اولیه شروع می‌شود.

$$X_{IJ} = \text{RAND} \times (UB - LB) + LB \quad K = 1, 2, \dots, N \quad (4)$$

که در آن X_{IK} ماتریس اولیه، $j = 1, 2, \dots, N$ نشان دهنده اندازه جمعیت (ردیف‌ها ماتریس اولیه) و n ابعاد (ستون‌های ماتریس اولیه) را نشان می‌دهد LB ، UB و rand حد پایین کران، بالا حد محدود و مقادیر تصادفی تولید شده را نشان می‌دهند. پروسه شکار

کروکودیل‌ها شامل دو فرآیند رفتارهای محاصره (اکتشاف) و شکار (استخراج) می‌باشد. در الگوریتم RSA کل فرآیند تکرار به چهار قسمت تبدیل می‌شود.

فاز اکتشاف (Exploration Phase)

در فاز اکتشاف یا محاصره کروکودیل اجازه نزدیک شدن به طعمه را ندارد هدف کروکودیل‌ها از فعالیت‌شان افزاز کردن کل فضای جستجو یا تضمین جستجوی سراسری (GLOBAL) است. فاز محاصره یا اکتشاف شامل پیاده‌روی به صورت شکمی یا ایستاده می‌شود که تضمین کننده جستجوی سراسری می‌باشد. یعنی این فاز عملاً می‌تواند نواحی متراکم را بعد از چندین بار تلاش پیدا کند. ۵۰٪ از فرآیند ابتدایی تکرار مربوط خواهد شد به فاز اکتشاف که ۲۵٪ ابتدای تکرار به پیاده‌روی ایستاده و ۲۵٪ تکرار دوم به پیاده‌روی شکمی اختصاص خواهد یافت.

$$X(t+1) = BESTK(t)(-\mu(t)) \times \beta - (R(t) \times RAND) \quad t \leq \frac{T}{4} \quad (3)$$

$$X(t+1) = BESTK(t) \times X(R1, K) \times ES(t) \times RAND \quad t \leq \frac{2T}{4} \text{ AND } t > \frac{T}{4} \quad (4)$$

که در آن $BESTK(t)$ راه‌حل بهینه به دست آمده در موقعیت k ام است $rand$ نشان دهنده عدد تصادفی، t عدد تکرار فعلی و حداکثر تعداد را نشان می‌دهد. تکرارها با T نشان داده می‌شود. $\mu(J, K)$ مقدار عملگر شکار جواب J ام در جایگاه k است و از معادله ۵ محاسبه می‌شود.

$$\mu(j, k) = bestk(t) \times p(j, k) \quad (5)$$

که در آن b یک پارامتر حساسیت است و دقت اکتشاف را توضیح می‌دهد. یکی دیگر تابعی به نام $R(J, K)$ که هدف آن کاهش فضای جستجو است، طبق معادله ۶ محاسبه می‌شود.

$$r(j, k) = bestk(t) - \frac{p(r2, k)}{bestk(t) + \epsilon} \quad (6)$$

که در آن $r1$ مقدار عدد تصادفی است که بین ۱ و N قرار دارد. در این جا N نشان دهنده است تعداد کل راه‌حل‌های نامزد $z(r1, l)$ یک موقعیت تصادفی برای k را نشان می‌دهد راه‌حل $r2$ نیز یک عدد دلخواه (تصادفی) بین ۱ و N است، در حالی که e نشان دهنده مقدار کمی است. $ES(t)$ که به عنوان یک حس تکاملی شناخته می‌شود، نسبت مثبتی بر احتمال حس تکاملی را می‌توان به صورت ریاضی نشان داد.

$$ES(t) = 2 \times r3 \times \left(1 - \frac{1}{T}\right) \quad (7)$$

که در آن $r3$ یک عدد تصادفی را نشان می‌دهد. $P(J, K)$ را می‌توان طبق معادله ۸ محاسبه کرد.

$$p(j, k) = \alpha + x(j, k) - \frac{m(xj)}{bestk(t) \times (ub(k) - lb(k)) + \epsilon} \quad (8)$$

که در آن a یک حد حساسیت است که دقت اکتشاف را کنترل می‌کند. $m(X, J)$ میانگین محل حل J ام است که می‌توان طبق معادله ۹ محاسبه کرد.

$$m(xj) = \frac{1}{n \sum_{k=1}^n x(j, k)} \quad (9)$$

فاز استخراج یا شکار (Exploitation Phase)

در فاز استخراج هدف از رویکرد همکاری و هماهنگی شکار نزدیک شدن به شکار یا تضمین جستجوی محلی (LOCAL) است. مرحله شکار نیز مانند مرحله محاصره دو راهبرد دارد، یعنی شکار هماهنگی و همکاری هر دوی این استراتژی‌ها برای عبور از فضای جستجو



استفاده می‌شوند مرحله شکار هم بر اساس تکرار به دو قسمت تقسیم می‌شود. استراتژی هماهنگی شکار و همکاری شکار ۲۵٪، سوم شکار مربوط به فرآیند هماهنگی شکار می‌شود که طبق معادله ۱۰ محاسبه می‌شود.

$$x(j, k)(t+1) = bestk(t) \times p(j, k(t)) \times rand \quad t \leq 3 \frac{T}{4} \text{ and } t > 2 \frac{T}{4} \quad (10)$$

و ۲۵٪ مرحله بهره‌برداری مربوط خواهد شد به نحوه همکاری برای شکار که محاسبه آن همانند معادله ۱۱ می‌باشد.

$$x(j, k)(t+1) = bestk(t) - \mu(j, k)(t) \times \varepsilon - r(j, k)(t) \times rand \quad t \leq T \text{ and } t > 3 \frac{T}{4} \quad (11)$$

که در آن $BESTK(t)$ موقعیت k در بهترین راه حل به دست آمده در تکرار جاری است. به طور مشابه، $\mu(J, K)$ عملگر شکار را نشان می‌دهد که با معادله ۵ محاسبه می‌شود انگیزه اصلی پشت الگوریتم RSA یافتن روش‌های جستجوی قدرتمندی است که می‌تواند راه‌حل‌هایی با کیفیت بهتر برای استخراج الگوهای معنادار و بهبود عملکرد پیش‌بینی ارائه دهد [۲۲، ۲۳].

در این پژوهش از یک شبکه عصبی عمیق پیش‌خور استفاده شد. این شبکه عصبی شامل سه لایه پنهان است. معماری این شبکه عصبی چهار لایه متراکم کلی دارد. سه لایه متراکم اول دارای تابع فعال‌سازی خطی و آخرین لایه متراکم دارای تابع فعال‌سازی سیگموئید است. معماری انتخاب شده امکان انعطاف‌پذیری در یادگیری الگوها و نمایش‌های پیچیده از داده‌های ورودی را فراهم می‌کند. پارامترهای مؤثر در الگوریتم RSA و شبکه عصبی در جدول ۲ آمده است. برای انتخاب هائپرپارامترها، از جست و جوی شبکه ای (Grid Search) بر روی دامنه‌های ممکن مانند تعداد نرون‌ها [۴، ۸، ۱۶، ۳۲]؛ Dropout [۰/۱، ۰/۲، ۰/۳، ۰/۵]؛ MaxIter [۲۰، ۵۰، ۱۰۰]؛ استفاده شد. معیار انتخاب، کمینه‌سازی خطای MSE و بیشینه‌سازی دقت در داده‌های اعتبارسنجی بود. این فرآیند با 10-fold cross-validation انجام شد تا از overfitting جلوگیری شود.

جدول ۲: هائپر پارامترهای الگوریتم RSA و شبکه عصبی در مدل پیشنهادی

شرح	مقدار	هائپر پارامترها
تعداد لایه‌های پنهان در شبکه عصبی پیش‌خور (FFN)	۳	تعداد لایه‌های پنهان
برای استخراج ویژگی‌های ابتدایی با تابع فعال‌سازی ReLU	۱۶	تعداد نرون‌ها در لایه پنهان یک
برای تحلیل عمیق‌تر ویژگی‌ها با تابع فعال‌سازی ReLU	۸	تعداد نرون‌ها در لایه پنهان دو
برای استخراج ویژگی‌های پیچیده با تابع فعال‌سازی ReLU	۶	تعداد نرون‌ها در لایه پنهان سه
برای طبقه‌بندی باینری (بیمار/سالم) با تابع فعال‌سازی Sigmoid	۱	تعداد نرون‌ها در لایه خروجی
برای یادگیری غیرخطی و استخراج ویژگی‌های پیچیده	(Rectified Linear Unit)	توابع فعال‌سازی لایه‌های پنهان
برای تولید خروجی باینری (۰ یا ۱)	سیگموئید (Sigmoid)	تابع فعال‌سازی لایه خروجی
بهینه‌سازی پویا با الگوریتم RSA برای کمینه‌سازی MSE	تنظیم‌شده توسط RSA	نرخ یادگیری
اعمال در لایه‌های پنهان برای جلوگیری از بیش‌برازش (Overfitting)	محدوده آزمایش‌شده: ۰/۲ تا ۰/۵ (۰/۳)	نرخ Dropout
تعداد تکرارهای آموزش شبکه عصبی	۵۰	تعداد دوره‌های آموزشی (EPOCHS)
برای بهینه‌سازی وزن‌ها و بایاس‌های شبکه عصبی	RSA (Reptile Search Algorithm)	الگوریتم بهینه‌سازی
برای همگرایی الگوریتم RSA	۵۰	تعداد تکرار RSA (MAXITER)
برای تعادل بین جست و جوی محلی و جهانی در RSA	۲۰	تعداد عامل‌های جست و جوی (SearchAgents)
وزن‌ها و بایاس‌های شبکه عصبی	محاسبه‌شده بر اساس تعداد نرون‌ها و لایه‌ها	تعداد پارامترهای مسئله DIM
بهینه‌شده توسط RSA برای هر لایه	$We1, We2, We3, bi1, bi2, bi3$	وزن‌ها و بایاس‌ها
داده‌های آموزشی و برچسب‌ها (بیمار/سالم)	TEX, TEY	پارامترهای ورودی شبکه
ارزیابی عملکرد شبکه بعد از هر آموزش	MSE	معیار خطا
توقف آموزش در صورت عدم بهبود عملکرد شبکه در مرحله آموزش	تا رسیدن به همگرایی یا تکرار ۵۰	شرط خاتمه

ساختار لایه‌های مدل پیشنهادی در شکل ۲ آورده شده است لایه ورودی این لایه شامل داده‌های ورودی است که ویژگی‌های مربوط به بیماران را دریافت می‌کند. لایه‌های پنهان: لایه پنهان ۱ دارای ۱۶ نورون با تابع فعال‌سازی ReLU برای استخراج ویژگی‌های ابتدایی لایه پنهان ۲ دارای ۸ نورون و همچنین با تابع فعال‌سازی ReLU برای تجزیه و تحلیل عمیق‌تر و لایه پنهان ۳ دارای ۶ نورون و با تابع فعال‌سازی ReLU که ویژگی‌های پیچیده‌تر را بررسی می‌کند لایه‌های استخراج ویژگی این لایه‌ها به‌طور کلی برای شناسایی و استخراج ویژگی‌های مهم از داده‌های ورودی استفاده می‌شوند. در این مورد، از لایه‌های پنهان (Hidden Layers) به عنوان لایه‌های استخراج ویژگی استفاده شده است. لایه خروجی این لایه تنها ۱ نورون دارد و از تابع فعال‌سازی Sigmoid برای تولید پیش‌بینی نهایی استفاده می‌کند، که در مواردی مانند پیش‌بینی بیماری کبدی که تنها یک نتیجه باینری (مثلاً وجود یا عدم وجود بیماری) را فراهم می‌آورد، از لایه خروجی جهت طبقه‌بندی پیش‌بینی نهایی استفاده شد.

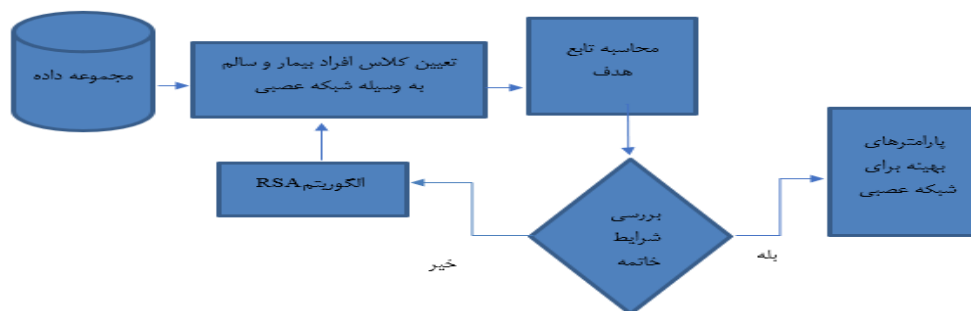


شکل ۲: ساختار شبکه عصبی عمیق مدل پیشنهادی

شبکه عصبی یکی از روش‌های طبقه‌بندی تقریبی برای پیش‌بینی است. برای آموزش وزن‌های شبکه عصبی، از یک شبکه عصبی با ساختار مشخص و معین، یعنی تعداد ثابتی لایه پنهانی و نورون در هر لایه برای شبکه موردنظر استفاده شد فرآیند آموزش وزن‌های شبکه با همان معماری مشخص آغاز شده و تا انتهای آموزش همان ساختار برای شبکه بود. به روزرسانی ساختار شبکه عصبی تبدیل به یک مسئله بهینه می‌شود. از جهتی الگوریتم جستجوی خزندگان، یک روش هوش جمعی و جدید برای یافتن نقاط بهینه سراسری در فضای جستجو است؛ بنابراین می‌توان از الگوریتم RSA برای پیدا کردن ساختار بهینه در شبکه عصبی استفاده کرد. در روش پیشنهادی برای پیش‌بینی بیماران کبدی، یک سیستم ترکیبی از شبکه عصبی و الگوریتم RSA استفاده شده است که در آن برای بالا بردن دقت و صحت سیستم و فرار از به دام افتادن الگوریتم در مینیمم محلی، ابتدا الگوریتم RSA وزن‌های بهینه شده را که به دقت قابل قبولی رسیده است، را به الگوریتم آموزشی پیشخور منتقل می‌کند. به عبارت دیگر، در روش پیشنهادی تمرکز بر این است که نقطه شروع شبکه عصبی، به گونه‌ای انتخاب شود که در مینیمم‌های محلی گیر نکنند. در الگوریتم ترکیبی RSA و FNN، در مراحل اولیه، پس از به دست آوردن بهترین ساختار شبکه عصبی توسط الگوریتم RSA، وزن‌های شبکه عصبی نیز توسط این الگوریتم پیدا می‌شوند. پس از یک تعداد تکرار معین یا پس از رسیدن به یک دقت از پیش تعیین شده، الگوریتم یادگیری از الگوریتم RSA به الگوریتم FNN سوییچ می‌شود. به این ترتیب الگوریتم پس انتشار خطا با یک شرایط اولیه خوب و در نزدیکی نقطه بهینه سراسری شروع می‌شود و طی تکرارهای کم به دقت بالایی می‌رسد؛ بنابراین در مدل پیشنهادی، علاوه بر بهینه کردن وزن‌های شبکه عصبی، ساختار آن نیز توسط الگوریتم RSA بهینه می‌شود. در واقع یک الگوریتم RSA در بدنه سیستم پیشنهادی وجود دارد، که فرآیند بهینه‌سازی ساختار شبکه عصبی را انجام می‌دهد. یک جمعیت از کروکودیل‌ها در داخل این بدنه وجود دارد که فرآیند بهینه‌سازی وزن‌های شبکه عصبی را انجام می‌دهد. همچنین به طور همزمان از الگوریتم RSA و تکنیک Dropout برای بهینه‌سازی وزن‌های شبکه عصبی استفاده شد. این دو روش به طور همزمان توانستند به بهبود کارایی و تعمیم‌پذیری مدل کمک کنند. به این صورت در حالی که شبکه عصبی تحت آموزش است، Dropout با غیرفعال کردن تصادفی نورون‌ها در هر دوره به جلوگیری از Overfitting کمک می‌کند. این کار باعث می‌شود که مدل به یادگیری ویژگی‌های مشترک و محاوره‌ای بپردازد و از یادگیری وابسته به یک مجموعه خاص از ویژگی‌ها جلوگیری کند. از Dropout در لایه‌های پنهان استفاده شد و Dropout با مقدار (۰/۲ تا ۰/۵) پیکربندی شد تا از Overfitting جلوگیری شود. با توجه به داده‌ها و جهت کنترل overfitting می‌توان با حذف لایه Fully connected و با استفاده از لایه‌های پنهان و تعداد کمتر نورون‌ها پیچیدگی کنترل شود

و به عملکرد بهتری دست یافت. استفاده از لایه پنهان می‌تواند باعث سادگی مدل و تسهیل در آموزش شود. همچنین از الگوریتم فرا ابتکاری برای بهینه‌سازی وزن‌های شبکه استفاده شد. این الگوریتم طی چند دور بهینه‌سازی به تدریج بهترین وزن‌ها را پیدا نمود. بعد از آموزش با استفاده از Dropout و بهینه‌سازی با الگوریتم فرا ابتکاری، عملکرد مدل روی مجموعه داده‌های اعتبارسنجی به صورت مرتب ارزیابی شد. پارامترهای Dropout و الگوریتم فرا ابتکاری به طور تجربی تنظیم شدند.

مجموعه داده به دو قسمت مجموعه آزمون و آموزش تبدیل شد که ۳۰ درصد داده‌ها به منظور آزمون و ۷۰ درصد آن به منظور آموزش استفاده شد. علت تقسیم داده‌ها این است که شبکه عصبی فقط با یکسری داده‌های خاص آموزش نبیند. اگر داده‌های جدیدی به عنوان ورودی به آن داده شود طبق معیارهای ارزیابی دقت، صحت و امتیاز fl عملکرد قابل قبولی داشته باشد. سپس الگوریتم بهینه ساز RSA فراخوانی می‌شود، و در هر مرحله وزن‌های شبکه عصبی را بهینه می‌کند. داده‌های آموزشی و تعداد نوره‌ها به ورودی الگوریتم RSA داده می‌شود. به وسیله تابع هدف الگوریتم RSA، وزن‌ها، بایاس‌ها و خطای شبکه عصبی محاسبه می‌شود این مراحل به شبکه عصبی کمک می‌کند تا به مرور و با تکرارهای مختلف، به دقت بالاتری در پیش‌بینی برسد. در مرحله آموزش در صورتی که پس از ۵۰ بار تکرار بهبود خطا بهبود نیافت آموزش متوقف خواهد شد. مراحل آموزش شبکه عصبی توسط الگوریتم RSA در شکل ۳ نشان داده شده است این شکل نمایانگر مراحل آموزش شبکه عصبی برای پیش‌بینی می‌باشد.



شکل ۳: آموزش شبکه عصبی در مدل پیشنهادی

در مدل پیشنهادی، بعد از ایجاد ساختار شبکه عصبی و مشخص کردن وزن‌ها، بایاس‌ها و الگوریتم آموزشی میانگین مربعات MSE پس انتشار خطا به کمک تابع خطا هزینه محاسبه می‌شود. این تابع مقدار موردنظر آموزش را از خروجی شبکه عصبی کم می‌کند و بهترین حالت این است که خروجی تابع MSE صفر شود.

$$MSE = \frac{1}{N} \sum (YI - YI') \quad (12)$$

که در آن N تعداد داده‌ها، YI داده‌های واقعی، YI' داده‌های تخمین زده شده می‌باشد. پس از انجام مرحله آموزش شبکه عصبی در مدل پیشنهادی، وارد مرحله آزمون و آزمایش مدل پیشنهادی جهت سنجش عملکرد شبکه عصبی در مدل پیشنهادی می‌شود. نحوه انجام مرحله آزمون شبکه عصبی در شکل ۴ نشان داده شده است. پس از انجام آموزش مدل پیشنهادی، یکسری پارامتر از مدل آموزش به دست می‌آید. به عنوان ورودی در مرحله آزمون از آن‌ها استفاده می‌شود علاوه بر پارامترهای $we1, we2, we3, bi1, bi2, bi3$ پارامتر x برای آزمون نیز وجود دارد که اکنون باید شبکه عصبی با پارامترهای بهینه شده آزمایش مورد پیش‌بینی قرار دهد.



شکل ۴: آزمایش شبکه عصبی در مدل پیشنهادی

کارایی مدل پیشنهادی با استفاده از معیارهای ارزیابی، به کمک ماتریس درهم‌ریختگی در جدول ۳ نشان داده شده است. اگر تعلق به دسته افراد بیمار را مثبت بودن (Positive) و عدم تعلق به این دسته را منفی بودن (Negative) در نظر گرفته شود. هر نمونه یا فردی در واقعیت، متعلق به یکی از کلاس‌های مثبت یا منفی است و از سوی دیگر، از هر الگوریتمی که برای دسته‌بندی داده‌ها استفاده شود، در نهایت هر نمونه عضو یکی از این دو دسته (Class) دسته‌بندی خواهد شد. بنابراین برای هر نمونه داده، یکی از چهار حالتی که در ادامه بیان شده، ممکن است اتفاق بیفتد.

نمونه عضو دسته مثبت باشد و عضو همین کلاس تشخیص داده شود مثبت صحیح یا TP (True Positive)

نمونه عضو کلاس مثبت باشد و عضو کلاس منفی تشخیص داده شود منفی کاذب یا FN (False Negative)

نمونه عضو کلاس منفی باشد و عضو همین کلاس تشخیص داده شود منفی صحیح یا TN (True Negative)

و در نهایت، نمونه عضو کلاس منفی باشد و عضو کلاس مثبت تشخیص داده شود مثبت کاذب یا FP (False Positive)

جدول ۳: ماتریس درهم‌ریختگی مدل

	پیش‌بینی مثبت (بیمار)	پیش‌بینی منفی (سالم)
مثبت واقعی (بیمار)	۳۹۹	۱۷
منفی واقعی (سالم)	۵	۱۶۲

در این مطالعه کارایی مدل پیشنهادی با استفاده از معیارهایی مانند، موارد فراخوانی (۱۲)، دقت (۱۳) صحت (۱۴) اندازه‌گیری شده است.

$$RECALL = \frac{TP}{TP + TN} \quad (12)$$

$$PRECISION = \frac{TP}{TP + FP} \quad (13)$$

$$ACCURACY = \frac{TP + TN}{TP + TN + FP + FN} \quad (14)$$

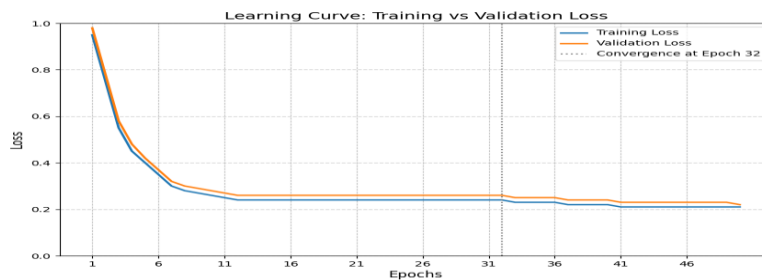
نتایج

در این مطالعه مدل پیشنهادی برای مجموعه داده ILPD موجود در مخزن یادگیری ماشین UCI ارزیابی شد. بدین منظور، مدل بر روی یک کامپیوتر محلی با ۷ هسته پردازشی، که شامل چندین ماشین مجازی با ۱۶ گیگابایت رم است، با استفاده از نرم‌افزار پایتون پیاده‌سازی شد. به منظور کاهش انحراف تصادفی، هر یک از الگوریتم‌ها ۵۰ بار اجرا شدند. داده‌ها جهت تشخیص افراد بیمار کبدی و افراد سالم مورد بررسی قرار گرفت و دقت تشخیص مدل مورد ارزیابی قرار گرفت، نتایج در جدول ۴ آورده شده است.

جدول ۴: نتایج اعمال مدل پیشنهادی

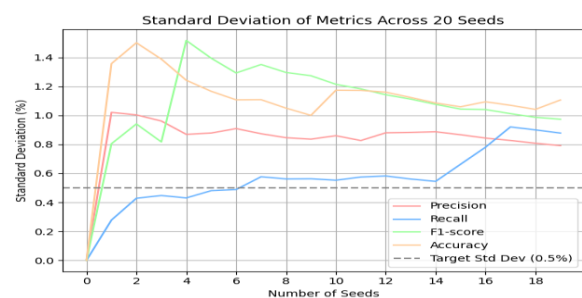
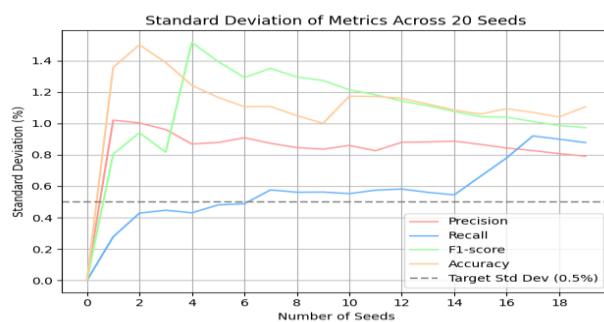
	تشخیص کلاس‌بندی	تشخیص کلاس‌بندی
	Class weighting	k-fold cross validation
Precision	۹۶/۹۸	۹۷/۱۲
Recall	۹۶/۰۶	۹۶/۹
F1-score	۹۶/۹۳	۹۷/۶
Accuracy	۹۶/۱۹	۹۶/۲۳

منحنی یادگیری مدل پیشنهادی در نمودار ۱ نشان داده شده است. این نمودار خطای آموزش (Training Loss) و خطای اعتبارسنجی (Validation Loss) را در طول ۵۰ مرحله آموزش نشان می‌دهد. خطاها از ابتدا کاهش می‌یابند و از epoch 32 به بعد همگرا می‌شوند، که بیانگر تعمیم‌پذیری خوب مدل و عدم وجود بیش‌برازش است. محور افقی تعداد epochها و محور عمودی مقدار خطا را نشان می‌دهد.



نمودار ۱: منحنی یادگیری مدل پیشنهادی

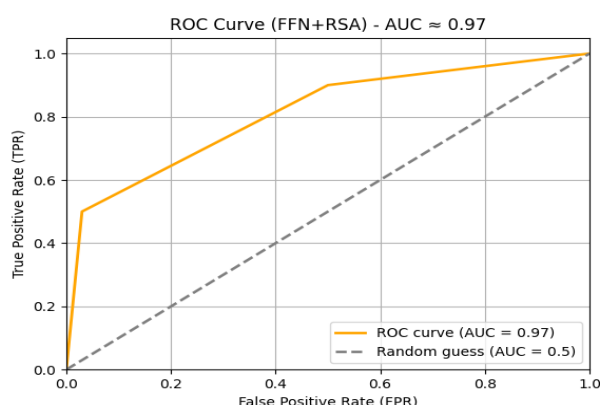
نمودارهای ۲ و ۳ به ترتیب پراکندگی انحراف معیار، معیارهای ارزیابی مدل‌های وزن‌دهی کلاس (Class Weighting) و اعتبارسنجی متقاطع k-تایی (k-fold cross-validation) را نشان می‌دهند. این نمودارها، پراکندگی انحراف معیارهای Recall، Precision، F1-score و Accuracy را در ۲۰ تکرار مستقل با استفاده از seedهای تصادفی (از ۱ تا ۲۰) به تصویر می‌کشند. مقادیر میانگین اولیه برای این معیارها در مدل وزن‌دهی کلاس به ترتیب ۹۶/۹۸٪، ۹۶/۰۶٪، ۹۶/۹۳٪ و ۹۶/۱۹٪ و در مدل اعتبارسنجی متقاطع k-تایی به ترتیب ۹۷/۱۲٪، ۹۶/۹٪، ۹۷/۶٪ و ۹۶/۲۳٪ در نظر گرفته شده‌اند. پراکندگی حول این مقادیر با انحراف معیار تقریبی ۱٪ محاسبه شده است. محور افقی تعداد seedها (از ۰ تا ۲۰) و محور عمودی انحراف معیار (به درصد) را نشان می‌دهد. یک خط افقی در سطح ۰/۵٪ به عنوان هدف انحراف معیار برای پایداری آماری اضافه شده است. در ابتدا، انحراف معیارها بالاتر از این سطح هستند، اما با افزایش تعداد seedها، به این هدف نزدیک می‌شوند. این نمودارها ثبات آماری مدل در برابر تغییرات تصادفی را تأیید می‌کنند و نشان می‌دهند که مدل پیشنهادی از پایداری نسبی برخوردار است.



نمودار ۲: انحراف معیارهای ارزیابی مدل class weighting در ۲۰ seed نمودار ۳: انحراف معیارهای ارزیابی مدل k-fold cross-validation در ۲۰ seed

نتایج نشان می‌دهد که روش پیشنهادی با استفاده از شبکه عصبی عمیق و استفاده از بهینه‌ساز RSA و روش k-fold cross-validation جهت حل مشکل داده‌های نامتوازن، به میزان معیار دقت ۹۶/۹٪، صحت ۹۷/۱۲٪، یادآوری ۹۶/۹٪ و امتیاز f1 ۹۷/۶٪ دست یافته است که نسبت به روش class weighting در مرحله آزمایش در تشخیص بیماری عملکرد بهتری داشته است. نمودار ۴ منحنی ROC مدل پیشنهادی (FFN+RSA) را نشان می‌دهد، که بر اساس خروجی‌های لایه Sigmoid و داده‌های آزمون محاسبه شده است. این نمودار نرخ مثبت صحیح (TPR) را در مقابل نرخ مثبت کاذب (FPR) برای آستانه‌های مختلف تصمیم‌گیری

نمایش می‌دهد. مقدار $AUC \approx 0.97$ نشان‌دهنده توانایی بالای مدل در تمایز بین کلاس‌های بیمار و سالم است. این نتایج با ماتریس درهم‌ریختگی جدول ۳، $TP=399$ ، $FP=5$ ، $TN=162$ ، $FN=17$ سازگار هستند، که نرخ $TPR \approx 0.9591$ و $FPR \approx 0.0299$ را برای آستانه پیش‌فرض (0.5) نشان می‌دهد. AUC بالا تأیید می‌کند که دقت $96/23\%$ مدل ناشی از عملکرد قوی آن است و نه بیش‌برازش.



نمودار ۴: ROC مدل پیشنهادی

جهت عادلانه بودن شرایط اجرا مدل‌های پیشین بر روی سیستم با مشخصات نام برده عملکرد این مطالعه با هشت روش مختلف SVM, LOGISTIC, K-NEARESTBORS, CATBOOST+HHO, CATBOOST+WOA, MLP, در جدول ۵ معیارهای F1-Score, Precision, Recall و Accuracy برای هر روش محاسبه و ارائه شده است.

جدول ۵: مقایسه مدل پیشنهادی با مدل‌های پیشین

امتیاز f1	نرخ یادآوری	درستی	دقت	طبقه‌بند
۷۰/۷۴	۷۹/۶۷	۵۰/۹۲	۵۰/۹۶	Logistic [۱۱]
۸۴/۲۵	۷۱/۰۹	۷۳/۶۲	۷۴/۰۳	k-nearest neighbors [۱۴]
۸۸/۸۸	۶۷/۷۵	۶۹/۴۴	۷۰/۱۹	Random forest [۱۰]
۶۴/۸۱	۵۸/۳۳	۵۷/۴۰	۵۷/۶۹	SVM [۱۴]
۹۲/۲۹	۸۰/۶۴	۸۴/۲۹	۸۴/۶۱	MLP [۱]
۹۶/۶	۹۶/۰۹	۹۷/۱۲	۹۶/۲۳	ffn+RSA (مدل پیشنهادی)
-	۹۶/۲۵	۹۶/۲۵	۹۵/۴۹	Ffn+psa [۲۴]
-	-	-	۸۲	Catbost+hho [۲۵]
-	-	-	۸۵	Catboost+woa [۲۵]

همان‌طور که در جدول ۵ مشاهده می‌شود، الگوریتم FFN-RSA در مقایسه با روش‌های مطرح‌شده در جدول از دقت بالاتری برخوردار است. یکی از دلایل این بهبود می‌تواند به توانایی الگوریتم RSA در یافتن مقادیر بهینه برای وزن‌های شبکه عصبی FFN مربوط باشد. RSA با ماهیت جستجوی سراسری در فضای فرایارامترها قادر است وزن‌های مناسبی برای شبکه عصبی بیابد که این خود منجر به کاهش خطای آموزش و بهبود دقت می‌شود. همچنین، ترکیب RSA با FFN موجب می‌شود که FFN به‌طور همزمان از مزایای جستجوی سراسری RSA و قابلیت یادگیری عمیق شبکه عصبی بهره‌مند شود، که این ترکیب به بهبود معیارهایی مانند Precision و F1-score منجر می‌شود. علاوه بر این، بهبود F1-score در مقایسه با سایر روش‌ها ناشی از توانایی ترکیب FFN-RSA در کاهش همزمان نرخ خطاهای False Positive و False Negative است، که این امر منجر به بهبود تعادل بین دقت (Precision) و یادآوری (Recall) می‌شود. به این ترتیب، ترکیب RSA با شبکه عصبی باعث می‌شود که مدل به تشخیص دقیق‌تر دسته‌ها با نرخ خطای کمتری دست یابد. روش ترکیبی پیشنهادی به دقت برابر با $96/23\%$ رسید که در مقایسه با هفت روش قبلی به $90/86\%$ بهبود



نسبت به مدل FFN-PSO رسید که می‌توان دلیل این بهبود را در تنظیمات وزن و بهینه سازی آن در لایه‌های پنهان با استفاده از الگوریتم فراابتکاری و Dropout دانست و از نظر درستی به مقدار ۹۷/۱۲٪ رسید که نسبت به سایر روش‌ها مقدار بهینه‌تری می‌باشد و با بهبودی برابر با ۱/۱۳٪ می‌توان ادعا کرد که روش پیشنهادی کلاس‌های بیشتری را درست تخمین زده است. ترکیب RSA-FFN در نرخ یادآوری به مقدار ۹۶/۰۹٪ رسید که از سایر روش‌های مورد مقایسه مقدار بالاتری است؛ اما نسبت به روش FFN-PSO، ۱۶٪ کمتر است. در مورد امتیاز F1 مدل به دقت ۹۶/۶٪ رسید که نسبت به سایر مدل‌های مقایسه شده ۴/۲۱٪ بهبود داشته است.

بحث و نتیجه گیری

نتایج این مطالعه نشان‌دهنده برتری مدل پیشنهادی مبتنی بر ترکیب شبکه عصبی پیشخور (FFN) و الگوریتم بهینه‌سازی جستجوی خزندگان (RSA) در پیش‌بینی بیماری کبدی با استفاده از مجموعه داده ILPD است. این مدل با دستیابی به دقت ۹۶/۲۳٪، نرخ یادآوری ۹۶/۹٪، دقت ۹۷/۱۲٪ و امتیاز F1 برابر با ۹۶/۱۲٪ در روش k-fold cross-validation، عملکردی برجسته ارائه می‌دهد که از بسیاری از روش‌های موجود پیشی گرفته است. مدل پیشنهادی در برابر تکنیک هیبریدی SMOTE-ENN با مدل‌های ensemble مانند AdaBoost و KNN که به دقت ۹۳/۲٪ رسیده‌اند [۱۹]، با بهره‌گیری از بهینه‌سازی عمیق فراپارامترها توسط RSA، به دقت بالاتری دست یافته است. این برتری به دلیل توانایی RSA در جستجوی مؤثر فضای فراپارامترها و ترکیب آن با تکنیک Dropout برای پیشگیری از بیش‌برازش است که تعمیم‌پذیری مدل را تقویت می‌کند. مدل‌های مبتنی بر boosting مانند Cat Boost و الگوریتم‌های HHO (Harris Hawks Optimization) و WOA (Whale Optimization Algorithm) به ترتیب دقت‌های ۸۲٪ و ۸۵٪ گزارش کرده‌اند [۲۵]، اما مدل پیشنهادی به دلیل بهره‌برداری از جستجوی ترکیبی سراسری و محلی RSA، عملکرد بهتری ارائه می‌دهد. در پژوهش انجام شده توسط آنان [۲۴]، ترکیب FFN با PSO به دقت ۹۵/۴۹٪ رسیده، مدل پیشنهادی ما بهبود ۰/۸۶٪ در دقت را نشان می‌دهد، که به جستجوی کارآمدتر RSA و فرار از نقاط محلی بهینه نسبت داده می‌شود. رویکرد مبتنی بر یادگیری projection و ranking نیز با دقت ۹۰/۱۲٪ [۲۳]، نرخ یادآوری کمتری (نسبت به ۹۶/۹٪ مدل مطالعه حاضر) ارائه کرده است، که اهمیت آن در تشخیص زودهنگام بیماری کبدی برجسته است. روش‌های کلاسیک ماشین لرنینگ مانند لجستیک رگرسیون با دقت ۷۵٪ [۲۰] یا مدل‌های ensemble stacking با دقت‌های ۹۳/۸۸٪ تا ۹۴/۱۲٪ [۹]، در مقایسه با مدل پیشنهادی، کارایی پایین‌تری دارند. پایداری مدل این مطالعه، که از طریق k-fold cross-validation (K=۱۰) و مدیریت عدم تعادل کلاس‌ها با وزن‌دهی معکوس حاصل شده، از بیش‌برازش جلوگیری کرده و تعمیم‌پذیری بالایی را تضمین می‌کند (انحراف معیار معیارهای عملکرد زیر ۰/۵٪). این مدل قابلیت ادغام در سیستم‌های تصمیم‌یار بالینی (CDSS) و اپلیکیشن‌های سلامت دیجیتال را دارد تا غربالگری زودهنگام را در جمعیت‌های پرخطر تسهیل کند، مرگ‌ومیر را کاهش دهد و هزینه‌های درمانی را تا ۵۰٪ کم کند [۷]. در نهایت، ادغام الگوریتم‌های فراابتکاری مانند RSA با شبکه‌های عصبی می‌تواند معیار جدیدی برای پیش‌بینی بیماری‌های کبدی ایجاد کند و برای مجموعه داده‌های مشابه در مطالعات آینده قابل تعمیم باشد برای تحقیقات آینده، پیشنهاد می‌شود مدل با شبکه‌های پیشرفته‌تر مانند CNN یا RNN ترکیب شود تا ویژگی‌های زمانی یا تصویری (مانند تصاویر سونوگرافی) را پوشش دهد. همچنین، افزایش حجم داده‌ها با ادغام مجموعه‌های بزرگ‌تر (مانند BUPA) و ترکیب RSA با الگوریتم‌های مبتنی بر گرادین می‌تواند عملکرد را بهبود بخشد. در نهایت، ارزیابی مدل در محیط‌های بالینی واقعی با داده‌های زمان واقعی، گامی کلیدی برای کاربردی‌سازی آن خواهد بود. در نتیجه، این پژوهش با ارائه یک مدل بهینه‌شده، گام مهمی در جهت تشخیص دقیق و زودهنگام بیماری کبدی برداشت و پتانسیل ادغام هوش مصنوعی در مراقبت‌های بهداشتی را برجسته کرد.

تعارض منافع

ندارد.

حمایت مالی

ندارد.

کد اخلاق

این مطالعه با کد IR.KMU.REC.1404.045 در کمیته اخلاق دانشگاه علوم پزشکی کرمان مورد تأیید قرار گرفته است.

سهم مشارکت نویسندگان

تمامی نویسندگان در این پژوهش سهم یکسان همکاری داشته اند.

References

- [1]. Ray R. Advances in data mining: Healthcare applications. International Research Journal of Engineering and Technology (IRJET) 2018;5(03):2395-0056.
- [2]. Kefelegn S, Kamat P. Prediction and analysis of liver disorder diseases by using data mining technique: survey. International Journal of Pure and Applied Mathematics 2018;118(9):765-70.
- [3]. Angulo P. Nonalcoholic fatty liver disease. N Engl J Med 2002;346:1221-31. doi: 10.1056/NEJMra011775.
- [4]. Shaheamlung, G, Kaur H, Kaur M. A Survey on machine learning techniques for the diagnosis of liver disease. International Conference on Intelligent Engineering and Management (ICIEM); 2020 Jun; 17-19; London, UK: IEEE; 2020. doi:10.1109/ICIEM48762.2020.9160097
- [5]. Durai V, Ramesh S, Kalthireddy D. Liver disease prediction using machine learning. Int J Adv Res Ideas Innov Technol 2019;5(2):1584-8.
- [6]. Tanwar N, Rahman KF. Machine Learning in liver disease diagnosis: Current progress and future opportunities. IOP Conference Series Materials Science and Engineering 2021; 1022(1): 12029. doi:10.1088/1757-899X/1022/1/012029
- [7]. UK Parliament. Women's health outcomes: Is there a gender gap? House of Lords Library; 2021 [cited 2024 Jan 8]. Available from: <https://lordslibrary.parliament.uk/womens-health-outcomes-is-there-a-gender-gap/>
- [8]. Jin H, Kim S, Kim J. Decision factors on effective liver patient data prediction. International Journal of Bio-Science and Bio-Technology 2014;6(4):167-78. <http://dx.doi.org/10.14257/ijbsbt.2014.6.4.16>
- [9]. Singh AS, Irfan M, Chowdhury A. Prediction of Liver Disease using Classification Algorithms. 4th International Conference on Computing Communication and Automation (ICCCA); 2018 Dec; 14-15; Greater Noida, India: IEEE; 2018. doi:10.1109/CCAA.2018.8777655
- [10]. Pathan A, Mhaske D, Jadhav S, Bhondave R, Rajeswari K. Comparative study of different classification algorithms on ILPD dataset to predict liver disorder. International Journal for Research in Applied Science & Engineering Technology (IJRASET) 2018;6(2):38-94.
- [11]. Adil SH, Ebrahim M, Raza K, Ali SS, Hashmani MA. Liver Patient Classification using Logistic Regression. 4th International Conference on Computer and Information Sciences (ICCOINS); 2018 Aug 13-14; Kuala Lumpur, Malaysia: IEEE; 2018. p. 1-5. doi: 10.1136/bmjhci-2021-100457
- [12]. Baitharu TR, Pani SK. Analysis of data mining techniques for healthcare decision support system using liver disorder dataset. Procedia Computer Science 2016;85:862-70. doi: 10.1016/J.PROCS.2016.05.276
- [13]. Geetha C, Arunachalam AR. Evaluation based approaches for liver disease prediction using machine learning algorithms. International Conference on Computer Communication and Informatics (ICCCI); 2021 Jan 27-29; Coimbatore, India: IEEE; 2021. p. 1-4.
- [14]. Vijayarani S, Dhayanand S. Liver disease prediction using SVM and Naïve Bayes algorithms. International Journal of Science, Engineering and Technology Research (IJSETR) 2015;4(4):816-20.
- [15]. Abdar M, Yen NY, Hung JC. Improving the diagnosis of liver disease using multilayer perceptron neural network and boosted decision trees. Journal of Medical and Biological Engineering 2018;38(6):953-65. doi:10.1007/s40846-017-0360-z
- [16]. Ramana BV, Boddu RK. Performance comparison of classification algorithms on medical datasets. 9th Annual Computing and Communication Workshop and Conference (CCWC); Las Vegas, NV, USA: IEEE; 2019. p. 140-5. doi: 10.1109/CCWC.2019.8666497
- [17]. Aswathy C. Liver patient dataset classification using the Intel® distribution for python. Intel, specialized development tools [cited 2022 Aug 12]. Available: <https://www.intel.com/content/www/us/en/developer/articles/technical/liver-patient-dataset-classification-usi>.
- [18]. Montazeri M, Montazeri M. Machine learning models for predicting the diagnosis of liver disease. Koomesh 2014;16(1):53-9. [In Persian]
- [19]. Khameneh ME, Sepehri MM, Saberifiroozi M. Using data mining for identify patients at high risk to hepatocellular carcinoma in the cirrhosis liver: preliminary report. Govareh 2014;19(4):265-74.
- [20]. Beale R, Jackson T. Neural Computing-an introduction. Translated by Alborzi M. Tehran: Sharif University of Technology Publications; 1990.
- [21]. Książek W, Abdar M, Acharya UR, Pławiak P. A novel machine learning approach for early detection of hepatocellular carcinoma patients. Cognitive Systems Research 2019;54:116-27.
- [22]. Khan MK, Zafar MH, Rashid S, Mansoor M, Moosavi SK, Sanfilippo F. Improved reptile search optimization algorithm: application on regression and classification problems. Appl Sci 2023;13(2):945. <https://doi.org/10.3390/app13020945>
- [23]. Garg R, Bhargava A. Bug prediction based on deep neural network with reptile search optimization to enhance software reliability. Multimedia Tools and Applications 2024;83(31):75869-91.
- [24]. Anand L, Syed Ibrahim SP. HANN: a hybrid model for liver syndrome classification by feature assortment optimization. Journal of Medical Systems 2018;42(11):211. doi: 10.1007/s10916-018-1073-8
- [25]. Mondal PK, Byeon H. Classification of Liver Disease Using Conventional Tree-Based Machine Learning Approaches with Feature Prioritization Using a Heuristic Algorithm. International Journal of Advanced Computer Science & Applications 2024;15(11). doi: 10.14569/IJACSA.2024.0151134