



Daily Streamflow Prediction of the Simineh River Using Machine and Deep Learning Models Enhanced with the Aquila Optimizer Algorithm

Edris Merufinia^{1*}, Kamran Yousefi², Mehrang Dusti Rezaei³ and Farzaneh Daroughe⁴

1-Department of civil engineering, SR.C., Islamic Azad University, Water Resources Expert, West Azerbaijan Regional Water Company, Urmia, Iran.

2-Surface Water Studies Expert, Basic Studies Office, West Azerbaijan Regional Water Company, Urmia, Iran

3-Manager of Basic Studies, West Azerbaijan Regional Water Company, Urmia, Iran

4-Water Resource Specialist, Maharab Omran Gostar Consulting Engineering Company, Iran

Received:2026-02-05

Accepted:2026-09-19

Revised:2026-07-18

Published online:2026-06-20

* Corresponding author email: edris.marufynia1389@gmail.com

ARTICLE INFO

ABSTRACT

Keywords:

River Discharge Prediction, Machine Learning, Random Forest, Aquila Optimization Algorithm, Simineh River

Introduction

Accurate prediction of river discharge is a fundamental requirement for effective water resources management, flood risk mitigation, reservoir operation, and climate change adaptation strategies. However, process-based hydrological models are often associated with considerable uncertainties arising from parameter estimation, model structure, and the complex nonlinear nature of rainfall-runoff processes, particularly in semi-arid and mountainous catchments. These challenges are exacerbated in regions with limited or heterogeneous hydro-meteorological observations. In recent years, data-driven approaches, including machine learning (ML) and deep learning (DL) models, have emerged as powerful alternatives to traditional physically based models due to their ability to capture nonlinear relationships without explicit representation of physical processes. Among various ML and DL techniques, Artificial Neural Networks (ANN), Random Forests (RF), and Long Short-Term Memory (LSTM) networks have been widely applied to hydrological forecasting problems. Nevertheless, the performance of these models strongly depends on the selection of appropriate input variables, hyperparameter tuning, and the balance between model complexity and generalization capability. Metaheuristic optimization algorithms have recently gained attention as efficient tools for automated hyperparameter optimization, enabling improved predictive accuracy and robustness. The Aquila Optimization (AO) algorithm, inspired by the hunting behavior of eagles, is a relatively new population-based metaheuristic that has demonstrated strong global search capability and fast convergence. The primary objective of this study is to evaluate and compare the performance of ANN, RF, and LSTM models optimized using the Aquila Optimization algorithm (ANN-AO, RF-AO, and LSTM-AO) for daily river discharge prediction. The Simineh River basin, located in West Azerbaijan Province, Iran, was selected as the study area due to its hydrological importance and exposure to hydroclimatic variability. Special emphasis is placed on assessing the impact of different input combinations and identifying the most robust and generalizable modeling scenario.

How to cite:

Edris Merufinia, E. Yousefi, K. Dusti Rezaei, M. Daroughe, F. *Daily Streamflow Prediction of the Simineh River Using Machine and Deep Learning Models Enhanced with the Aquila Optimizer Algorithm*. Journal of Hydraulics and Water Science, 36 (1):83-102
<https://doi.org/10.22034/hws.2026.72908.1043>



This is an open-access article under the CC BY NC license
(<https://creativecommons.org/licenses/by-nc/4.0/>)



Materials and Methods

Daily hydro-meteorological data spanning a 40-year period (1982–2022) were used in this study. Observed river discharge and precipitation data were obtained from the Iran Water Resources Management Company, while additional hydroclimatological variables were extracted from NASA POWER satellite products. The input variables included bias-corrected precipitation (mm/day), two-meter air temperature range (°C), relative humidity at two meters (%), wind speed range at ten meters (m/s), and antecedent river discharge with one- to three-day time lags (m³/s). The target output variable was daily river discharge (m³/s). Prior to modeling, an extensive data preprocessing procedure was conducted, including homogeneity testing, detection and removal of outliers, reconstruction of missing values, and linear normalization. Pearson correlation coefficient (PCC) analysis was employed to quantify the relationships between candidate input variables and river discharge. Based on this analysis, seven different input scenarios were defined to systematically investigate the role of hydroclimatic variables and hydrological memory in model performance. The dataset was divided into training (70%) and testing (30%) subsets. The Aquila Optimization algorithm was applied to optimize the hyperparameters of all three models to ensure a fair and objective comparison. Model performance was evaluated using Root Mean Square Error (RMSE), coefficient of determination (R^2), and Nash–Sutcliffe Efficiency (NSE) in both training and testing phases, with particular emphasis on testing performance as an indicator of generalization capability.

Results and Discussion

The results demonstrate that all three optimized models were capable of capturing the general dynamics of daily river discharge; however, their predictive accuracies and generalization behaviors varied considerably across different input scenarios. For the ANN-AO model, the best testing performance was achieved under Scenario SN6, which exhibited balanced errors between training and testing phases, indicating reduced overfitting. Nevertheless, the overall predictive accuracy of ANN-AO was lower than that of RF-AO and LSTM-AO. The RF-AO model consistently outperformed the other models in terms of robustness and generalization. In the testing phase, Scenario SN4 yielded the best results with RMSE = 11.931 m³/s, R^2 = 0.834, and NSE = 88.0%. This scenario combined key hydroclimatic variables with selected discharge lags, achieving an optimal balance between model complexity and predictive power. Notably, scenarios with excessively rich input combinations, such as SN7, showed slightly better performance during training but experienced degradation in testing accuracy, highlighting the negative impact of over-parameterization on generalization. The LSTM-AO model demonstrated competitive performance, particularly in scenarios with limited but informative inputs. The best testing performance for LSTM-AO was observed under Scenario SN1 (RMSE = 10.541 m³/s, R^2 = 0.862, NSE = 86.0%). However, when all hydroclimatic variables and discharge lags were included, LSTM performance deteriorated, suggesting increased sensitivity to noise and reduced robustness under daily time-step conditions. Overall, the comparative analysis revealed that RF-AO provided the most stable and accurate predictions across scenarios, outperforming LSTM-AO and ANN-AO by approximately 2.3% and 36.0% in NSE, respectively. These findings emphasize that increasing the number of input variables does not necessarily improve predictive performance and that intelligent feature selection plays a more critical role than model complexity alone.

Conclusion

This study confirms the high potential of data-driven models, particularly Random Forest combined with metaheuristic optimization, for reliable daily river discharge prediction in hydroclimatically complex basins. The superior performance of RF-AO highlights the effectiveness of ensemble-based algorithms in reducing variance and enhancing generalization without requiring explicit sequential memory or complex network architectures. The results further demonstrate that excessive input enrichment can negatively affect model robustness, especially in ANN and LSTM frameworks. The findings provide valuable insights for developing intelligent decision-support systems for water resources management, flood forecasting, and climate change impact assessment in data-scarce regions.

Acknowledgement

The authors gratefully acknowledge the Iran Water Resources Management Company for providing hydrological data and NASA POWER for supplying hydroclimatological datasets used in this study.



نشریه دانش آب و هیدرولیک

شاپا الکترونیکی: 3092-6114
درگاه نشریه: <https://hws.tabrizu.ac.ir/>



مقاله پژوهشی

پیش بینی روزانه جریان رودخانه سیمینه با استفاده از مدل های یادگیری ماشین و عمیق تقویت شده با مدل بهینه سازی عقاب

ادریس معروفی نیا^{۱*}، کامران یوسفی^۲، مهرنگ دوستی رضایی^۳ و فرزانه داروغه^۴

۱- دانش آموخته دکتری عمران آب دانشگاه آزاد اسلامی واحد علوم و تحقیقات تهران، کارشناس منابع آب شرکت منطقه ای آذربایجان غربی، ارومیه، ایران

۲- کارشناس مطالعات آب های سطحی دفتر مطالعات پایه شرکت آب منطقه ای آذربایجان غربی، ارومیه، ایران

۳- مدیر مطالعات پایه شرکت آب منطقه ای آذربایجان غربی، ارومیه، ایران

۴- کارشناس منابع آب شرکت مهندسی مشاور مهارآب عمران گستر، ایران

تاریخ پذیرش: ۱۴۰۵/۰۶/۲۸

تاریخ دریافت: ۱۴۰۴/۱۱/۱۶

تاریخ انتشار آنلاین: ۱۴۰۵/۰۳/۳۱

تاریخ ویرایش: ۱۴۰۵/۰۴/۲۷

نویسنده مسئول: edris.marufynia1389@gmail.com

چکیده

پیش بینی جریان رودخانه بر اساس فرایندها با عدم قطعیت های قابل توجهی در پارامترهای مدل و روش های پارامتری سازی مرتبط با فرایندهای پیچیده تولید رواناب مواجه است. مدل های داده محور جایگزین هایی کارآمد ارائه می دهند که نیازی به در نظر گرفتن فرایندهای فیزیکی ندارند. آمار و اطلاعات بارش و دبی جریان ۴۰ ساله رودخانه سیمینه رود در استان آذربایجان غربی در مقیاس روزانه از سال ۱۹۸۲ تا ۲۰۲۲ از شرکت مدیریت منابع آب ایران و اطلاعات هیدروکلیماتولوژی از ماهواره پاورلارس ناسا دریافت شد. جهت پیش بینی جریان رودخانه از مدل های یادگیری ماشین شبکه عصبی مصنوعی مدل پرسپترون چندلایه (MLP)، مدل جنگل تصادفی (RF) و مدل یادگیری عمیق شبکه حافظه طولانی کوتاه مدت (LSTM) استفاده شد. برای بهینه نمودن ضرایب و ابرپارامترهای مدل نیز از الگوریتم بهینه سازی (فرایتنکاری) عقاب (AO) استفاده شد. عملیات پیش پردازش داده ها (از قبیل آزمون همگنی، نرمال سازی داده و حذف داده های پرت و بازسازی اطلاعات مفقوده) قبل از فرآیند مدل سازی انجام گرفت. پارامترهای ورودی تحقیق شامل بارش تصحیح شده (میلی متر بر روز)، دامنه تغییرات دمای دو متری (درجه سانتی گراد)، رطوبت نسبی در ارتفاع دو متری (درصد)، دامنه تغییرات سرعت باد در ارتفاع ده متری (متر بر ثانیه)، دبی جریان با یک روز تا سه روز تأخیر (مترمکعب بر ثانیه) و دبی جریان رودخانه (مترمکعب بر ثانیه) به عنوان خروجی استفاده شد. برای انتخاب سناریو و ترکیب مدل از ضریب همبستگی پیرسون (PCC) بین متغیرهای ورودی و پارامتر خروجی استفاده شد و مجموعاً از هفت سناریو استفاده شد. تفکیک و تقسیم بندی داده ها برای فرآیند آموزش و آزمون با نسبت ۷۰ به ۳۰ انجام گرفت. نتایج نشان داد که مدل RF-AO در سناریوی چهارم (SN4) با ترکیب بهینه ای از متغیرهای هیدروکلیماتیک، بهترین تعمیم پذیری را در فاز آزمون داشته است؛ به طوری که مقادیر $RMSE = 11.931$ ، $R^2 = 0.834$ و $NSE = 88\%$ را کسب کرد. این مدل نسبت به LSTM-AO و ANN-AO به ترتیب حدود ۲.۳٪ و ۳۶.۰٪ بهبود در معیار NSE نشان داد. یافته ها همچنین تأکید می کنند که افزودن بیش از حد ورودی ها، عملکرد را بهبود نمی بخشد و انتخاب هوشمندانه ترکیب ورودی ها عامل کلیدی در دستیابی به پیش بینی های قابل اعتماد است. برتری RF-AO، ظرفیت بالای روش های داده محور ساده تر را در شبیه سازی سیستم های هیدرولوژیک پیچیده برجسته می سازد و زمینه را برای توسعه سامانه های تصمیم یار هوشمند در مدیریت آب فراهم می آورد.

کلمات کلیدی: الگوریتم عقاب، پیش بینی جریان رودخانه، ضریب پیرسون، رودخانه سیمینه، مدل یادگیری ماشین

۱- مقدمه

پیش‌بینی جریان رودخانه‌ها در حوضه‌های آبریز برای کنترل سیلاب، مدیریت منابع آب و سایر مسائل مهندسی ضروری است. از یک سو، شدت سیلاب‌های جهانی توسعه پایدار اجتماعی-اقتصادی را تهدید می‌کند و خسارات گسترده‌ای را به دشت‌های سیلابی یا مناطق با جمعیت بالا در پایین‌دست رودخانه‌های اصلی وارد می‌سازد (Wang et al., 2024). پیش‌بینی جریان رودخانه‌ها می‌تواند اطلاعات هشدار سریع به‌هنگام و دقیق در مورد سیلاب‌ها ارائه دهد، روند و میزان سیلاب‌ها و تغییرات سطح آب را فراهم کند و همچنین یک پایه علمی برای پیشگیری و واکنش به سیلاب‌ها ارائه دهد (Xu et al., 2024). از سوی دیگر، پیش‌بینی جریان رودخانه‌ها سنگ بنای تخصیص بهینه منابع را تشکیل می‌دهد و یک پایه دقیق برای تصمیم‌گیری در توزیع و مدیریت منابع آب ارائه می‌دهد تا به این ترتیب کارایی بهره‌برداری از منابع آب افزایش یافته و توسعه پایدار این منبع ارزشمند تضمین شود (Shi et al., 2023). مدل‌های پیش‌بینی جریان رودخانه‌ها به مدل‌های فیزیکی و مدل‌های مبتنی بر داده تقسیم‌بندی می‌شوند. مدل‌های فیزیکی معمولاً بر مفاهیم بنیادی هیدرولوژیکی استوار هستند تا تکامل رودخانه، رواناب و فرآیندهای تولید جریان را شبیه‌سازی کنند (Yao et al., 2023). مدل‌های یادگیری ماشین این قابلیت را دارند که از طریق یک فرآیند مدل‌سازی ساده، روابط پیچیده بین ویژگی‌ها را درک کرده و عوامل تأثیرگذار را به‌طور دقیق تجزیه و تحلیل کنند، بدون آنکه نیاز به درک عمیق از فرآیندهای فیزیکی زمینه‌ای داشته باشند (Lin et al., 2023). استفاده از مدل‌های یادگیری ماشین (هوش مصنوعی) مختلف مبتنی بر داده از قبیل رگرسیون بردار پشتیبان، جنگل‌های تصادفی، تقویت گرادیان شدید، شبکه عصبی بیزین، شبکه عصبی مصنوعی، سیستم‌های فازی و استنتاج عصبی-فازی تطبیقی و... برای پیش‌بینی جریان رودخانه طی دهه گذشته به‌شدت مورد پژوهش قرار گرفته است (Chen et al., 2024). با این حال، یک رابطه غیرخطی پیچیده بین متغیرهای پیش‌بین و متغیرهای هدف وجود دارد که اغلب منجر به بیش‌برازش در هنگام اعمال تکنیک‌های کلاسیک یادگیری ماشینی بر روی مجموعه داده‌های بزرگ می‌شود (Zhang & Yan, 2023). مدل‌های یادگیری عمیق، گسترشی از یادگیری ماشین هستند که به روی هم‌گذاری سطوح عمیق مفاهیم اشاره دارند تا رایانه‌ها قادر به یادگیری مفاهیم پیچیده شوند و از توانایی بالایی در انجام وظایف طبقه‌بندی و پیش‌بینی برخوردارند (Ng et al., 2023). مدل‌های یادگیری عمیق از لایه‌های متعدد برای استخراج ویژگی استفاده می‌کنند که آن‌ها را قادر می‌سازد تا به‌طور مؤثری روابط پیچیده در متغیرهای پیش‌بین را درک کنند (Guo et al., 2023). تعداد فزاینده‌ای از مدل‌های نوآورانه، نتایج بهبودیافته‌ای را در پیش‌بینی جریان رودخانه‌ها به دست آورده‌اند، با مثال‌هایی مانند شبکه حافظه طولانی-کوتاه مدت (LSTM)، شبکه‌های عصبی عمیق (DNNs) و

ترانسفورمر (TR)، شبکه کانولوشنی زمانی (TCN) و واحد بازگشتی دروازه‌ای (GRU) به‌طور متداول و گسترده در مسائل پیش‌بینی جریان رودخانه‌ها استفاده می‌شوند (Guo et al., 2023; Qiao et al., 2023). مانگوکیا و شاراما (۲۰۲۵) یک رویکرد مبتنی بر یادگیری عمیق برای ارتقای پیش‌بینی جریان رودخانه‌ها در حوضه‌های آبریز با مشاهدات تجمعی و متناوب را مورد ارزیابی قرار دادند. یافته‌های آن‌ها نشان داد که داده‌های فصول مرطوب برای بهبود پیش‌بینی‌های مدل بسیار مهم هستند و داده‌های هفتگی می‌توانند نتایجی قابل مقایسه با مشاهدات روزانه ارائه دهند (Mangukiya & Sharma, 2025). در تحقیق دیگر دانش و همکاران (۲۰۲۵) به ارزیابی مقایسه‌ای مدل‌های یادگیری ماشین و یادگیری عمیق برای پیش‌بینی جریان روزانه رودخانه پرداختند. در پژوهش مذکور از سه مدل یادگیری ماشین، جنگل تصادفی (RF)، درخت تصمیم (DT) و نزدیک‌ترین همسایگی (KNN) و سه چارچوب یادگیری عمیق شامل شبکه‌های عصبی کانولوشنی (CNN)، حافظه کوتاه‌مدت بلند مدت (LSTM) و یک مدل ترکیبی CNN-LSTM استفاده نمودند. یافته‌های آن‌ها نشان داد که مدل‌های یادگیری عمیق به‌طور مداوم از روش‌های یادگیری ماشین برای پیش‌بینی جریان رودخانه در هر دو مکان بهتر عمل می‌کنند. مدل ترکیبی CNN-LSTM دقیق‌ترین پیش‌بینی‌ها را ارائه داد (Danesh et al., 2025). در تحقیق دیگر، هوانگ و همکاران (۲۰۲۵) به ارزیابی مدل پیش‌بینی جریان روزانه مبتنی بر یادگیری عمیق برای حوضه رودخانه هانجیانگ پرداختند. در پژوهش مذکور از مدل‌های پیشنهادی، خودتوجهی (SA)، شبکه عصبی کانولوشنی یک‌بعدی (1D-CNN) و حافظه کوتاه‌مدت بلندمدت دوطرفه (BiLSTM) استفاده نمودند. یافته‌های آن‌ها نشان داد که مدل SA-CNN-BiLSTM پیشنهادی، به‌ویژه برای زمان‌های انتظار طولانی و رویدادهای سیل، دقت بهبود یافته قابل‌توجهی را نشان می‌دهد و کمی‌سازی عدم قطعیت قوی را فراهم می‌کند و در نتیجه ابزاری قابل‌اعتمادتر برای بهره‌برداری از مخزن و مدیریت ریسک سیل ارائه می‌دهد (Huang et al., 2025). در سال‌های اخیر، مدل‌های یادگیری ماشین و عمیق برای پیش‌بینی جریان رودخانه گسترش یافته‌اند، اما بسیاری از آن‌ها به‌ویژه شبکه‌های عصبی در مدل‌سازی نوسانات غیرخطی شدید (مانند سیلاب و خشکسالی) با مشکلات همگرایی به بهینه محلی و حساسیت به هاپرپارامترها مواجه‌اند. اگرچه برخی مطالعات از الگوریتم‌های فراابتکاری برای بهینه‌سازی استفاده کرده‌اند، ترکیب این الگوریتم‌ها با معماری‌های پیچیده DL مانند LSTM کمتر مورد توجه قرار گرفته و ارزیابی جامع آن‌ها در شرایط هیدرولوژیکی دشوار و با داده‌های محدود، هنوز ناکافی است. در این راستا، هدف اصلی این پژوهش، توسعه و مقایسه مدل‌های ترکیبی نوین بر پایه یادگیری عمیق کوپل‌شده با الگوریتم‌های فراابتکاری برای پیش‌بینی دقیق‌تر جریان رودخانه است. سایر اهداف تحقیق شامل بهینه‌سازی خودکار ابرپارامترها و ساختار مدل‌های پایه با حداقل وابستگی به

۲-۲- ساختار کلی پژوهش

در این پژوهش جهت پیش‌بینی جریان رودخانه سیمینه‌رود، در گام نخست به بررسی منطقه مورد مطالعه پرداخته شد. در گام بعدی به جمع‌آوری اطلاعات و آمار مورد نیاز جهت پیش‌بینی جریان (آمار هیدرومتری و هیدروکلیماتولوژی) پرداخته شد. جهت پیش‌بینی جریان و تخمین دبی خروجی جریان رودخانه از مدل‌های یادگیری ماشین شبکه عصبی مصنوعی (ANN)، جنگل تصادفی (RF) و مدل یادگیری عمیق شبکه عصبی حافظه طولانی کوتاه‌مدت (LSTM) استفاده شد. همچنین از مدل‌های بهینه‌سازی فراابتکاری دسته عقاب ها (AO) با شبکه عصبی ترکیب شده (ANN-AO) و مدل جنگل تصادفی نیز با الگوریتم مذکور ترکیب شده (RF-AO) و مدل حافظه بلندمدت کوتاه‌مدت نیز با الگوریتم مذکور ترکیب (LSTM-AO) شد. این ترکیب و هیبریدسازی مدل‌های بهینه‌سازی فراابتکاری با مدل‌های یادگیری ماشین و عمیق به منظور غلبه بر ضعف‌های ذاتی این مدل‌ها (از جمله گیر کردن در بهینه‌های محلی، همگرایی کند و حساسیت به پارامترهای اولیه) انجام گرفته است. الگوریتم‌های فراابتکاری با جستجوی هوشمند در فضای پارامتری، وزن‌ها و ساختار مدل را به گونه‌ای بهینه می‌کنند که دقت، پایداری و تعمیم‌پذیری آن را بهبود می‌بخشند. این ترکیب به‌ویژه در مسائل پیچیده و غیرخطی مانند مدل‌سازی جریان رودخانه که داده‌ها محدود و دارای نویز هستند، اهمیت دوچندانی دارد. برای پیش‌بینی و مدل‌سازی از برنامه MATLAB 2023b و برای ترسیم نمودار و گراف‌های مربوطه از برنامه‌های EdrawMax 15.0.2.1409 و OriginLab 2022 استفاده شد. قبل از شروع فرایند مدل‌سازی، عملیات پیش‌پردازش داده‌ها (از قبیل شناسایی و حذف داده‌های پرت، بی‌مقیاس‌سازی داده‌ها از طریق عملیات نرمال‌سازی و انجام آزمون همگنی) انجام شد. در گام بعد تقسیم‌بندی داده‌ها در فرایند آموزش و آزمون انجام گرفت. لذا از نسبت ۷۰ به ۳۰ درصد در فرایند مدل‌سازی برای آموزش و آزمون استفاده شد. داده‌ها به دو روش انتخاب متوالی و انتخاب تصادفی انتخاب می‌شوند. در این پژوهش، برای پیشگیری از کم‌برازش و بیش‌برازش، از روش انتخاب تصادفی استفاده شد تا روند آموزش مدل به نحو مطلوب‌تری انجام گیرد. پس از انجام پیش‌بینی و اتمام فرایند مدل‌سازی، ارزیابی نتایج با استفاده از معیارهای ارزیابی انجام گرفت. در نهایت به بحث و تحلیل نتایج و انتخاب مدل‌های برتر و جمع‌بندی پرداخته شد. شکل (۱) نمودار کلی فرایند انجام کار را نشان می‌دهد.

به منظور شناخت ویژگی‌های اقلیمی منطقه مورد مطالعه و ارزیابی اولیه کیفیت داده‌ها، پارامترهای آمار توصیفی متغیرهای هواشناسی شامل مقادیر بیشینه، کمینه، میانگین، انحراف استاندارد، چولگی و کشیدگی برای کل دوره آماری محاسبه شد. نتایج این تحلیل که در جدول (۱) ارائه شده است، نشان‌دهنده دامنه تغییرات گسترده در پارامترهای دمایی و بارشی است. بر اساس مقادیر چولگی و کشیدگی، توزیع فراوانی متغیرهایی نظیر سرعت باد و بارش دارای چولگی مثبت

تنظیم دستی، کاهش خطای پیش‌بینی در شرایط حدی (سیلاب و خشکسالی) و ارزیابی قابلیت تعمیم‌پذیری مدل‌های پیشنهادی در حوزه‌های آبریز با ویژگی‌های هیدرولوژیکی متفاوت است. نوآوری این پژوهش در تلفیق هم‌زمان مدل‌های یادگیری ماشین و یادگیری عمیق با الگوریتم فراابتکاری عقاب (AO) برای پیش‌بینی روزانه جریان رودخانه سیمینه است. رویکردی که با بهینه‌سازی خودکار ساختار مدل، ضرایب و اپرپارامترها، وابستگی به تنظیمات تجربی و دستی را کاهش می‌دهد. در این راستا، عملکرد سه مدل RF، MLP و LSTM در قالب نسخه‌های بهینه‌شده با الگوریتم AO به صورت مقایسه‌ای و در شرایط واقعی یک حوضه آبریز نیمه‌خشک ارزیابی شده است. همچنین، استفاده هم‌زمان از داده‌های هیدرولوژیکی زمینی و متغیرهای هیدروکلیماتولوژیکی ماهواره‌ای پاورلارس ناسا در مقیاس روزانه و طی یک دوره ۴۰ ساله، بستر مناسبی برای مدل‌سازی پایدار و قابل‌اعتماد فراهم کرده است. از سوی دیگر، انتخاب ترکیب بهینه ورودی‌ها بر پایه ضریب همبستگی پیرسون و طراحی هفت سناریوی ورودی، امکان بررسی نظام‌مند اثر متغیرها بر دقت پیش‌بینی را فراهم ساخته و اهمیت انتخاب هوشمندانه ورودی‌ها را، فراتر از صرف افزایش تعداد آن‌ها، نشان می‌دهد.

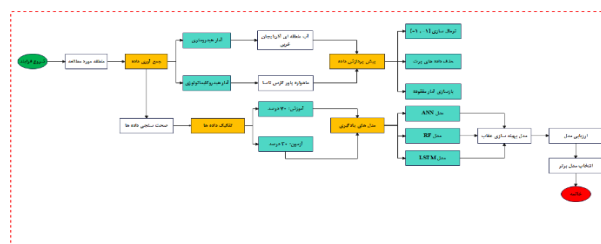
۲- مواد و روش‌ها

۲-۱- روش تحقیق و جمع‌آوری داده‌ها

در این پژوهش جهت پیش‌بینی جریان رودخانه سیمینه‌رود در استان آذربایجان غربی، ابتدا داده‌های بارش و دبی جریان رودخانه در مقیاس روزانه از سال ۱۹۸۲ تا ۲۰۲۲ از شرکت مدیریت منابع آب ایران دریافت شد. همچنین آمار دمای کمینه روزانه در ارتفاع ۲ متری از سطح زمین حسب درجه سانتی‌گراد (T2M_MIN)، دمای بیشینه روزانه در ارتفاع ۲ متری از سطح زمین حسب درجه سانتی‌گراد (T2M_MAX)، بارش تصحیح‌شده بر حسب میلی‌متر بر روز (PRECTOTCORR)، سرعت باد در ۱۰ متری حسب متر بر ثانیه (WS10M)، رطوبت نسبی در ۲ متری حسب درصد (RH2M)، فشار سطحی جو حسب کیلوپاسکال (PS)، بیشینه سرعت باد در ۱۰ متری حسب متر بر ثانیه (WS10M_MAX)، کمینه سرعت باد در ۱۰ متری حسب متر بر ثانیه (WS10M_MIN)، دامنه (تفاضل بین حداکثر و حداقل) سرعت باد روزانه در ارتفاع ۱۰ متری از سطح زمین حسب متر بر ثانیه (WS10M_RANGE)، جهت باد در ۱۰ متری حسب درجه (WD10M)، رطوبت مخصوص (ویژه) در ۲ متری حسب گرم بخار آب در هر کیلوگرم هوا (QV2M)، دامنه تغییرات دمای روزانه (یعنی اختلاف بین حداکثر و حداقل دما) در ارتفاع استاندارد ۲ متری از سطح زمین حسب درجه سانتی‌گراد (T2M_RANGE)، دما در ۲ متری حسب درجه سانتی‌گراد (T2M) و سرعت باد در ۲ متری حسب متر بر ثانیه (WS2M) از ماهواره پاورلارس ناسا از سال ۱۹۸۰ تا سال ۲۰۲۰ در مقیاس روزانه دریافت شد.

ترکیب از ویژگی‌ها را برای فرآیند آموزش مدل توجیه می‌کند بر اساس نتایج تحلیل همبستگی پیرسون و به منظور ارزیابی تأثیر ترکیب‌های مختلف متغیرهای ورودی بر پیش بینی دبی خروجی جریان، هفت سناریوی مدل سازی طراحی شد که مشخصات آنها در جدول (۳) آمده است. در این سناریوها، متغیر خروجی در تمامی حالات، دبی گام زمانی جاری $Q(t)$ است. سناریوهای اول تا سوم یعنی SN_1 تا SN_3 صرفاً بر اساس مقادیر تاخیری دبی با گام های زمانی یک تا سه روز گذشته تعریف شده اند تا توانایی مدل در بازسازی خودهمبستگی جریان ارزیابی شود. در سناریوهای بعدی، متغیرهای اقلیمی بر اساس اولویت همبستگی آنها با دبی به تدریج به ساختار ورودی افزوده شده‌اند. بر این اساس، در سناریوی چهارم (SN_4) رطوبت نسبی، در سناریوی پنجم (SN_5) دامنه تغییرات دما، در سناریوی ششم (SN_6) بارندگی اصلاح شده و در نهایت در سناریوی هفتم (SN_7) دامنه تغییرات سرعت باد به متغیرهای ورودی اضافه شدند تا میزان اثرگذاری هر یک از این پارامترهای محیطی بر بهبود عملکرد مدل مشخص گردد

بوده که بیانگر رخدادهای حدی در این پارامترها است، در حالی که متغیرهای دمایی توزیعی نزدیک‌تر به حالت نرمال را نشان می‌دهند. این تحلیل آماری مبنای واسنجی و پیش‌پردازش داده‌ها در مراحل بعدی تحقیق قرار گرفته است. به منظور شناسایی مؤثرترین متغیرهای ورودی بر مدل سازی جریان، تحلیل همبستگی پیرسون بین متغیرهای هواشناسی، مقادیر تاخیری دبی و خروجی مدل (Q_t) انجام گرفت که نتایج آن در جدول (۲) ارائه شده است. بر اساس این نتایج، بیشترین میزان همبستگی مربوط به مقادیر تاخیری دبی است، به طوری که متغیر Q_{t-1} با ضریب همبستگی ۰.۹۱۵ قوی‌ترین ارتباط مستقیم را با دبی لحظه‌ای نشان می‌دهد. این امر بیانگر حافظه کوتاه‌مدت قوی در سیستم هیدرولوژیکی منطقه است. در میان متغیرهای هواشناسی، رطوبت نسبی ($RH2M$) با ضریب ۰.۲۴۶ و دامنه تغییرات دمایی ($T2M_RANGE$) با ضریب ۰.۱۷۱ بیشترین تأثیر را بر تغییرات دبی داشته‌اند. ضرایب مثبت در تمامی پارامترها حاکی از وجود رابطه مستقیم بین ورودی‌های منتخب و خروجی مدل است که انتخاب این



شکل ۱- نمودار کلی روند انجام تحقیق

Fig 1. Overall flowchart of the research methodology

جدول ۱- خلاصه پارامترهای آماری محدوده مطالعاتی برای کل داده‌ها

Table 1. Summary of statistical parameters for the study area based on the entire dataset

ردیف	پارامتر	بیشینه	کمینه	میانگین	انحراف استاندارد	چولگی	گشیدگی
۱	T2M_MIN	23.46	-22.42	5.107	8.532	-0.135	-1.016
۲	T2M_MAX	38.27	-9.48	16.534	11.016	-0.022	-1.240
۳	PRECTOTCORR	54.45	0.01	11.981	15.523	0.762	-1.223
۴	RH2M	98.81	14.19	56.773	18.558	-0.030	-0.978
۵	WS10M	12.96	0.71	3.313	1.554	1.187	1.915
۶	PS	83.81	80.81	82.578	0.386	-0.277	0.605
۷	WS10M_MAX	17.34	1.2	5.356	2.518	1.011	0.857
۸	WS10M_MIN	10.21	0.01	1.595	1.114	1.005	1.620
۹	WS10M_RANGE	13.59	0.51	3.761	1.916	1.138	1.236
۱۰	WD10M	345.69	10.56	182.553	71.639	-0.903	-0.549
۱۱	QV2M	12.02	0.85	5.092	2.053	0.433	-0.471
۱۲	T2M_RANGE	20.91	1.17	11.427	3.281	-0.129	-0.692
۱۳	T2M	30.17	-15.7	10.366	9.894	-0.037	-1.221
۱۴	WS2M	10.32	0.52	2.418	1.183	1.341	2.531
۱۵	Discharge	350.45	0.003	15.063	28.858	4.092	22.785

جدول ۲- ارتباط متغیرهای ورودی و خروجی تحقیق با استفاده از ضریب همبستگی پیرسون

Table 2. Correlation between input and output variables of the study using Pearson correlation coefficient

T2M_RANGE	RH2M	WS10M_RANGE	PRECTOTCORR	Q_{t-1}	Q_{t-2}	Q_{t-3}	Q_t
0.171	0.246	0.129	0.138	0.915	0.812	0.745	1

جدول ۳- ترکیب مدل و سناریوهای پژوهش بر اساس ضریب پیرسون

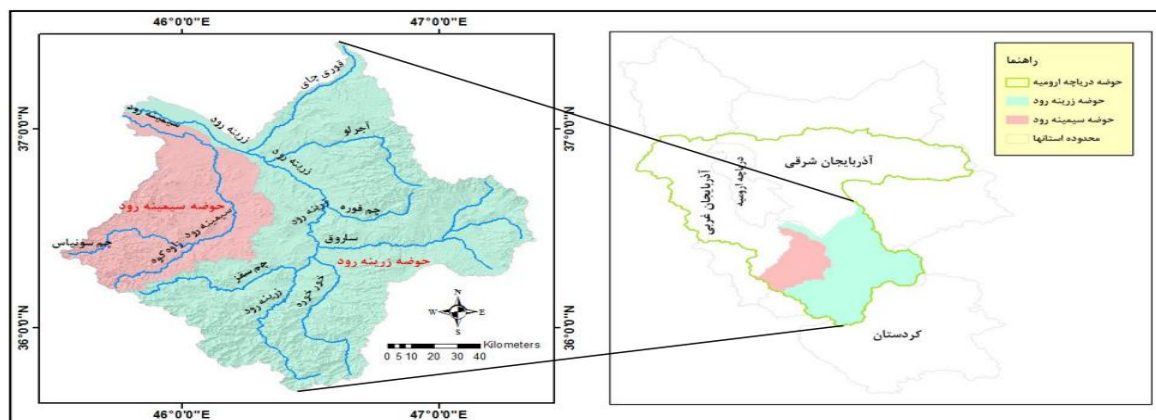
Table 3. Model combinations and research scenarios based on Pearson correlation coefficient

متغیر خروجی	ترکیب ورودی	سناریو
$Q(t)$	Q_{t-1}	SN ₁
$Q(t)$	Q_{t-1}, Q_{t-2}	SN ₂
$Q(t)$	$Q_{t-1}, Q_{t-2}, Q_{t-3}$	SN ₃
$Q(t)$	$Q_{t-1}, Q_{t-2}, Q_{t-3}, RH2M$	SN ₄
$Q(t)$	$Q_{t-1}, Q_{t-2}, Q_{t-3}, RH2M, T2M_RANGE$	SN ₅
$Q(t)$	$Q_{t-1}, Q_{t-2}, Q_{t-3}, RH2M, T2M_RANGE, PRECTOTCORR$	SN ₆
$Q(t)$	$Q_{t-1}, Q_{t-2}, Q_{t-3}, RH2M, T2M_RANGE, PRECTOTCORR, WS10M_RANGE$	SN ₇

کیلومتر مربع است و ارتفاع آن از ۱۲۸۶ متر تا حداکثر ۲۸۰۴ متر از سطح متوسط دریا متغیر است. همچنین، شیب متوسط حوضه ۱/۴٪ بوده و جهت کلی آن از جنوب غرب به سمت شمال شرق است. این حوضه به دلیل اهمیت آن برای محیط زیست دریاچه ارومیه و همچنین به دلیل وجود ایستگاه‌های پایش کافی برای مطالعات میدانی و دسترسی به داده‌ها، به عنوان منطقه مورد مطالعه انتخاب شد (Bostanmaneshrad et al., 2018). شکل (۲) منطقه مورد مطالعه این رودخانه را نشان می‌دهد.

۲-۳- منطقه مورد مطالعه

حوضه آبخیز سیمینه در شمال غرب ایران واقع شده است و بین طول جغرافیایی ۳۵°۴۵' تا ۳۵°۴۶' شرقی و عرض جغرافیایی ۳۶°۱۰' تا ۳۷° شمالی قرار دارد. این رودخانه از کوه‌های شمال غرب ایران سرچشمه می‌گیرد، مسیری حدود ۲۰۰ کیلومتری را به سمت شمال از میان شهرها، اراضی کشاورزی و مناطق صنعتی طی کرده و در نهایت به دریاچه ارومیه می‌ریزد. مساحت حوضه آبخیز سیمینه حدود ۳۸۸۴



شکل ۲- موقعیت حوضه‌های آبریز زرينه رود و سيمينه رود (احمد آلی و همکاران، ۱۳۹۶)

Fig 2. Location of the Zarrineh River and Simineh River basins

گسترده‌تری هستند (Janiesch et al., 2021; Wrisberg et al., 2002). در ادامه مدل‌های یادگیری ماشین و روابط ریاضی آن‌ها ارائه شده‌اند.

۲-۴-۱- مدل شبکه عصبی مصنوعی

ماهیت تطبیقی، توانایی یادگیری، شناسایی آسان و عملکرد سریع شبکه‌های عصبی مصنوعی، آن‌ها را به یکی از مدل‌های پرکاربرد در یادگیری ماشین تبدیل کرده است (Vatanchi et al., 2023). این سیستم‌ها به عنوان یک روش پیشرفته پردازش داده شناخته می‌شوند

۲-۴-۲- مدل‌های یادگیری ماشین و یادگیری عمیق

مدل‌های یادگیری ماشین و یادگیری عمیق، روش‌هایی هوشمند برای استخراج الگوها از داده‌ها هستند. مدل‌های یادگیری عمیق با به کارگیری شبکه‌های عصبی چندلایه، قادر هستند ویژگی‌های پیچیده‌تر را از داده‌های غیرساختاریافته مانند تصاویر یا سری‌های زمانی به صورت خودکار یاد بگیرد. این مدل‌ها به ویژه در مسائلی با حجم داده‌ی بالا و روابط غیرخطی پیچیده، عملکرد برتری از خود نشان می‌دهند، هرچند نیازمند منابع محاسباتی بیشتر و داده‌های آموزشی

تا (۸) مقادیر حالت سلول و حالت پنهان را برای یک سلول LSTM (در گام زمانی t) نشان می‌دهند:

$$f_t = \sigma(W_{fx}x_t + W_{fh}h_{t-1} + b_f) \quad (۴)$$

$$i_t = \sigma(W_{ix}x_t + W_{ih}h_{t-1} + b_i) \quad (۵)$$

$$o_t = \sigma(W_{ox}x_t + W_{oh}h_{t-1} + b_o) \quad (۶)$$

$$c_t = f_t \otimes c_{t-1} + i_t \otimes \tanh(W_{cx}x_t + W_{ch}h_{t-1} + b_c) \quad (۷)$$

$$h_t = o_t \otimes \tanh(c_t) \quad (۸)$$

که در روابط بالا به ترتیب که W : نشان‌دهنده پارامترهای وزن، b : بایاس، σ : تابع سیگموئید و $\otimes$ ضرب عنصر به عنصر است. هر زیرنویس در W و b ، وزن و بایاس یک دروازه خاص را نشان می‌دهد. به عنوان مثال، W_{fx} : وزن ورودی x_t در دروازه f_t است. همچنین c : حالت سلول، h : حالت پنهان، x : ورودی و $\tanh$ تابع تانژانت هیپربولیک است. همچنین اولین دروازه، دروازه فراموشی f_t است که اطلاعات را از حالت سلول قبلی کنترل می‌کند. این دروازه تعیین می‌کند که چه مقدار از داده‌ها باید حفظ شده یا به مرحله بعدی منتقل شوند. دروازه ورودی، دومین دروازه است که میزان ورود اطلاعات جدید را تعیین می‌کند.

۲-۴-۴- مدل بهینه‌سازی عقاب (آکویلا)

الگوریتم بهینه‌ساز عقاب یا آکویلا (AO)، یک الگوریتم فراابتکاری مبتنی بر جمعیت است که توسط Abualigah et al. (2021) ارائه شده است و از رفتار اجتماعی عقاب در شکار الهام گرفته شده است. فرایند بهینه‌سازی با جمعیتی از جواب‌های نامزد (X) آغاز می‌شود، همان‌گونه که در معادله (۹) نشان داده شده است. این جمعیت به صورت تصادفی در بازه بین حد بالایی (UB) و حد پایینی (LB) مسئله تولید می‌شود. بهترین جواب به دست آمده تا آن لحظه، در هر تکرار به عنوان تقریبی از جواب بهینه در نظر گرفته می‌شود.

$$X = \begin{bmatrix} x_{1,1} & \dots & x_{1,j} & x_{1, \text{Dim}-1} & x_{1, \text{Dim}} \\ x_{2,1} & \dots & x_{2,j} & \dots & x_{2, \text{Dim}} \\ \dots & \dots & x_{i,j} & \dots & \dots \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ x_{N-1,1} & \dots & x_{N-1,j} & \dots & x_{N-1, \text{Dim}} \\ x_{N,1} & \dots & x_{N,j} & x_{N, \text{Dim}-1} & x_{N, \text{Dim}} \end{bmatrix} \quad (۹)$$

که در این رابطه به ترتیب X : مجموعه‌ای از جواب‌های کاندیدای فعلی را نشان می‌دهد که به صورت تصادفی با استفاده از معادله (۱۰) تولید شده‌اند، X_i : مقادیر تصمیم‌گیری (موقعیت‌ها) جواب i ام را نشان می‌دهد، N : تعداد کل جواب‌های کاندیدا (جمعیت) است و Dim نیز بعد مسئله را مشخص می‌کند.

که به خوبی قادرند با رویدادهای هیدرولوژیکی پیچیده و غیرخطی سازگار شوند (Jimmy et al., 2020). شبکه پرسپترون چندلایه (MLP) را می‌توان به صورت ریاضی به شرح زیر بیان کرد:

$$y = f\left(\sum_{i=1}^n w_i p_i + b\right) \quad (۱)$$

که در این رابطه به ترتیب w_i بردار وزن، p_i بردار ورودی (برای $i = 1, \dots, n$)، f تابع انتقال و y خروجی است.

۲-۴-۲- مدل جنگل تصادفی

مدل جنگل تصادفی (RF) که توسط بریمن (Breiman, 2001) معرفی شد، یک روش یادگیری ماشین گروهی برای مسائل دسته‌بندی و رگرسیون است. این مدل بر پایه مجموعه‌ای از درخت‌های تصمیم (CARTs) ساخته می‌شود که به عنوان دسته‌بندی‌های ترکیبی شناخته می‌شوند؛ این درخت‌ها به صورت تصادفی پیش‌بین‌ها و داده‌های واسنجی را برای هر درخت مستقل انتخاب کرده و سپس تمامی درخت‌ها را ترکیب می‌کنند تا پیش‌بینی‌هایی قابل اطمینان‌تر از متغیر وابسته ارائه دهند (Bargam et al., 2024). در مدل جنگل تصادفی (RF)، دو پارامتر کلیدی تعداد درخت‌های موجود در جنگل (n_{tree}) و تعداد پیش‌بین‌هایی که در هر گره ارزیابی می‌شوند (m_{try}) باید تعیین شوند که توسط معادلات (۲) و (۳) تعریف شده‌اند:

$$m_{\text{try}} = \log_2(M + 1) \quad (۲)$$

$$\hat{y}(x_i) = \frac{1}{K} \sum_{k=1}^K T_b(x_i) \quad (۳)$$

که در روابط بالا به ترتیب m_{try} : تعداد پیش‌بین‌های هر گره، K : تعداد درختان جنگل و T_b : نشان‌دهنده هر درخت است.

۲-۴-۳- مدل حافظه طولانی-کوتاه مدت (LSTM)

هوخرایتر و شمیدهور (۱۹۹۷) معماری شبکه LSTM را معرفی کردند که به عنوان یک پیشرفت نسبت به شبکه‌های عصبی بازگشتی (RNN) در نظر گرفته می‌شود (Hochreiter & Schmidhuber, 1997). یک شبکه LSTM معمولی از مجموعه‌ای از بلوک‌های حافظه تشکیل شده است که به صورت ساختاری شناخته شده به عنوان سلول‌های LSTM پیچیده می‌شوند. این بلوک‌های حافظه نقشی مشابه نورون‌های موجود در لایه‌های پنهان ایفا می‌کنند (Akiner et al., 2024). در شبکه LSTM، هر سلول حافظه از طریق دو مؤلفه کلیدی به یکدیگر متصل است. حالت پنهان (h_t) که به عنوان حافظه کوتاه مدت شناخته می‌شود و حالت سلول (c_t) که حافظه بلندمدت محسوب می‌شود. شبکه LSTM قادر است داده‌های حیاتی را برای مدت زمان طولانی‌تری ذخیره کند (Le et al., 2019). معادلات (۴)

$$\sigma = \left(\frac{\Gamma(1 + \beta) \times \sin\left(\frac{\pi\beta}{2}\right)}{\Gamma\left(\frac{1+\beta}{2}\right) \times \beta \times 2^{\left(\frac{\beta-1}{2}\right)}} \right) \quad (۱۶)$$

که در این روابط S: مقدار ثابتی است که برابر با ۰/۱ در نظر گرفته شده است و u و v اعداد تصادفی در بازه هستند، β : یک مقدار ثابت برابر با ۰/۵ است، و σ یکی از پارامترهای تابع توزیع پرواز لوی را نشان می‌دهد. گام بعدی بهره‌برداری گسترده است. این مکانیزم در واقع پرواز پایین همراه با حمله‌ای با شیب کند است که در آن عقاب به تدریج در فضای هدف فرو می‌رود و به طعمه نزدیک‌تر می‌شود تا به آن حمله کند. این رفتار عقاب نیز به‌دقت با استفاده از معادله (۱۷) بیان شده است:

$$X_3(t+1) = (X_{best}(t) - X_M(t)) \times \alpha - \text{rand} + ((UB - LB) \times \text{rand} + LB) \times \delta, \quad (۱۷)$$

که در این رابطه $X_3(t+1)$ جواب مربوط به تکرار بعدی t است که توسط سومین روش جستجو (X_3) تولید می‌شود. $X_{best}(t)$ مکان تقریبی طعمه را تا تکرار i ام (یعنی بهترین جواب به‌دست‌آمده) نشان می‌دهد و $X_M(t)$ میانگین جواب‌های فعلی را در تکرار t ام بیان می‌کند. rand یک مقدار تصادفی بین ۰ و ۱ است. پارامترهای تنظیم‌کننده بهره‌برداری و در این مقاله به مقدار کوچکی (۰/۱) ثابت در نظر گرفته شده‌اند. α و δ حد پایین و حد بالای مسئله مورد نظر را نشان می‌دهند. در نهایت، مرحله آخر روش حمله عقاب الهام گرفته شده است که به‌طور گسترده با نام حمله راه‌رفتن و گرفتن شناخته می‌شود که مرحله آخر، یعنی بهره‌برداری متمرکز، بر آن استوار است. این مفهوم نیز به‌صورت ریاضی در معادله (۱۸) بیان شده است:

$$X_4(t+1) = QF \times X_{best}(t) - (G_1 \times X(t) \times \text{rand}) - G_2 \times \text{Levy}(D) + \text{rand} \times G_1 \quad (۱۸)$$

که در این رابطه به ترتیب $X_4(t+1)$: جواب مربوط به تکرار بعدی t است که توسط چهارمین روش جستجو (X_4) تولید می‌شود. QF: یک تابع کیفیت است که برای ایجاد تعادل میان راهبردهای جستجو به‌کار می‌رود و با استفاده از معادله (۱۹) محاسبه می‌شود. G_1 : انواع حرکات الگوریتم بهینه‌ساز عقاب را نشان می‌دهد که در طی فرآیند تعقیب طعمه استفاده می‌شوند و با معادله (۲۰) تولید می‌گردند. G_2 : مقادیری کاهشی از ۲ تا ۰ را ارائه می‌دهد که شیب پرواز عقاب را در حین تعقیب طعمه از موقعیت اول (۱) تا آخرین موقعیت (t) نشان می‌دهد و با استفاده از معادله (۲۱) محاسبه می‌شود. $X(t)$ نیز جواب فعلی در تکرار t ام است.

$$QF(t) = t^{\frac{2 \times \text{rand} - 1}{(1-T)^2}} \quad (۱۹)$$

$$X_{ij} = \text{rand} \times (UB_j - LB_j) + LB_j, i = 1, 2, \dots, N_j = 1, 2, \dots, \text{Dim} \quad (۱۰)$$

که در این رابطه که در آن rand یک عدد تصادفی است، LB_j : حد پایینی بعد از j ام و UB_j : حد بالایی بعد از j ام مسئله مورد نظر را نشان می‌دهد. مرحله بعدی اکتشاف گسترده است که شامل پرواز بلند همراه با غوطه‌وری عمودی است. در این روش، الگوریتم هدف دارد فضای جستجو را از ارتفاع زیاد بررسی کند. این مفهوم به‌صورت ریاضی در معادله (۱۱) نشان داده شده است.

$$X_1(t+1) = X_{best}(t) \times \left(1 - \frac{t}{T}\right) + (X_M(t) - X_{best}(t)) \times \text{rand} \quad (۱۱)$$

که در آن $X_1(t+1)$ جواب مربوط به تکرار بعدی t است که توسط اولین روش جستجو (X_1) تولید می‌شود. $X_{best}(t)$ بهترین جواب به‌دست‌آمده تا تکرار t ام را نشان می‌دهد که مکان تقریبی طعمه را منعکس می‌کند. عبارت $\left(1 - \frac{t}{T}\right)$ برای کنترل گستره جستجو (اکتشاف) در طول تکرارها به‌کار می‌رود. $X_M(t)$ میانگین مکان جواب‌های فعلی را در تکرار t ام نشان می‌دهد که با استفاده از معادله (۱۲) محاسبه می‌شود. rand یک مقدار تصادفی بین ۰ و ۱ است. همچنین t و T به‌ترتیب نشان‌دهنده تکرار فعلی و حداکثر تعداد تکرارها هستند.

$$X_M(t) = \frac{1}{N} \sum_{i=1}^N X_i(t), \forall j = 1, 2, \dots, \text{Dim} \quad (۱۲)$$

مرحله بعدی اکتشاف متمرکز است که شامل پرواز کنترلی همراه با سرخوردن کوتاه است. در این روش، عقاب با چرخیدن دور طعمه هدف، زمین را برای حمله آماده می‌کند. چنین رفتاری از عقاب به‌طور دقیق با استفاده از معادله (۱۳) به‌دست می‌آید:

$$X_2(t+1) = X_{best}(t) \times \text{Levy}(D) + X_R(t) + (y - x) \times \text{rand} \quad (۱۳)$$

که در این رابطه $X_2(t+1)$ جواب مربوط به تکرار بعدی t است که توسط دومین روش جستجو (X_2) تولید می‌شود. D: فضای ابعادی مسئله و $\text{Levy}(D)$ تابع توزیع پرواز لوی است که با استفاده از معادله (۱۷) محاسبه می‌شود. $X_R(t)$ یک جواب تصادفی است که در بازه $[1, N]$ در تکرار i ام انتخاب می‌شود. x و y شکل مارپیچی را در فرایند جستجو نشان می‌دهند که با استفاده از معادلات (۱۴) تا (۱۶) بیان می‌شود.

$$\text{Levy}(D) = s \times \frac{u \times \sigma}{|v|^{\frac{1}{\beta}}} \quad (۱۴)$$

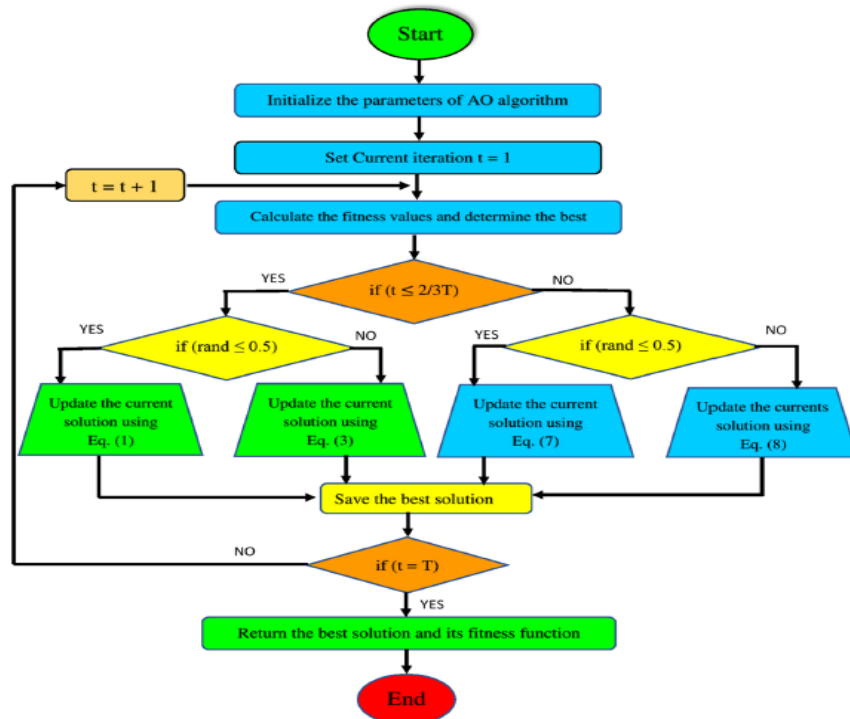
$$x = r \times \sin(\theta), y = r \times \cos(\theta) \quad (۱۵)$$

نشان‌دهنده تکرار فعلی و حداکثر تعداد تکرارها هستند. $Levy(D)$: تابع توزیع پرواز لوی است. شکل (۳) فلوچارت و مراحل کلی الگوریتم عقاب‌ها را نشان می‌دهد.

$$G_1 = 2 \times r \text{ and } -1 \quad (20)$$

$$G_2 = 2 \times \left(1 - \frac{t}{T}\right) \quad (21)$$

که در این روابط به ترتیب $QF(t)$: مقدار تابع کیفیت در تکرار t ام است و $rand$: یک عدد تصادفی بین ۰ و ۱ است. t و T به ترتیب



شکل ۳- فلوچارت روند کلی الگوریتم عقاب (Sasmal et al., 2023)

Fig 3. Flowchart of the overall procedure of the Aquila Optimizer algorithm

قوت و ضعف آن‌ها و اتخاذ تصمیمات آگاهانه در مدیریت منابع آب را فراهم می‌آورد. در این پژوهش از معیارهای ضریب تبیین (R^2)، میانگین قدرمطلق خطا (MAE)، ریشه میانگین مربعات خطا (RMSE)، بهره‌وری نش-ساتکلیف (NSE) استفاده شد. جدول (۴) معیارهای ارزیابی و محدوده مقادیر آنها را نشان می‌دهد.

۲-۵- معیارهای ارزیابی مدل

معیارهای ارزیابی نقشی اساسی در سنجش دقت، پایداری و قابلیت اطمینان مدل‌های پیش‌بینی جریان رودخانه ایفا می‌کنند. استفاده ترکیبی از چندین معیار برای اجتناب از قضاوت‌های یک‌سویه ضروری است. این رویکرد جامع، امکان مقایسه عادلانه مدل‌ها، شناسایی نقاط

جدول ۴- معیارهای ارزیابی مدل جهت پیش‌بینی جریان رودخانه

Table 4. Model evaluation metrics for streamflow prediction

ردیف	نام معیار	فرمول ریاضی	محدوده تعریف
۱	ضریب تبیین (R^2)	$R^2 = \frac{[\sum_{i=1}^n (Q_{obs,i} - \bar{Q}_{obs})(Q_{sim,i} - \bar{Q}_{sim})]^2}{\sum_{i=1}^n (Q_{obs,i} - \bar{Q}_{obs})^2 \cdot \sum_{i=1}^n (Q_{sim,i} - \bar{Q}_{sim})^2}$	$(-\infty, 1]$
۲	میانگین قدر مطلق خطا (MAE)	$MAE = \frac{1}{n} \sum_{i=1}^n Q_i^{obs} - Q_i^{sim} $	$[0, +\infty)$
۳	ریشه میانگین مربعات خطا (RMSE)	$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (Q_i^{obs} - Q_i^{sim})^2}$	$[0, +\infty)$
۴	ضریب بهره‌وری نش-ساتکلیف (NSE)	$NSE = 1 - \frac{\sum_{i=1}^n (Q_{sim,i} - Q_{obs,i})^2}{\sum_{i=1}^n (Q_{obs,i} - \bar{Q}_{obs})^2}$	$(-\infty, +\infty)$

دستی این پارامترها به دلیل ماهیت پیچیده و غیرخطی داده‌های منابع آب، اغلب منجر به قرارگیری در بهینه‌های محلی و کاهش دقت تعمیم‌دهی مدل می‌شود. در این پژوهش، برای دستیابی به ساختار بهینه و افزایش توان پیش‌بینی از الگوریتم بهینه‌سازی عقاب (AO) استفاده شد. جدول (۵) مقادیر ضرایب تنظیم مدل‌های پژوهش را نشان می‌دهد.

که در روابط بالا به ترتیب Q_i^{obs} : مقدار مشاهده‌شده (واقعی) در گام زمانی i ، Q_i^{sim} : مقدار شبیه‌سازی‌شده (پیش‌بینی‌شده) در گام زمانی i ، $\bar{Q}^{obs}$: میانگین مقادیر مشاهده‌شده و n نیز تعداد کل نمونه‌ها (داده‌های زمانی) است. واحد معیارهای خطای MAE و RMSE نیز مترمکعب در ثانیه است.

۲-۶- تنظیم ضرایب و ابرپارامترهای مدل‌های

پژوهش

عملکرد مدل‌های مبتنی بر یادگیری ماشین و یادگیری عمیق (MLP)، RF و LSTM) به شدت تابع انتخاب دقیق ابرپارامترها است. تنظیم

جدول ۵- تنظیم ضرایب و هایپرپارامترهای مدل پژوهش

Table 5. Hyperparameters Tuning and Model Coefficients Adjustment

ردیف	نام پارامتر	محدوده جستجو	مقدار بهینه	توصیف پارامتر
مدل پرسپترون چند لایه (MLP)				
۱	تعداد لایه‌های پنهان (Number of Hidden Layers)	۱-۳	۲	لایه‌های واسط که وظیفه استخراج ویژگی‌های غیرخطی از داده‌های ورودی را بر عهده دارند.
۲	تعداد نرون‌های لایه اول (Neurons in Hidden Layer 1)	۵-۱۰۰	۴۸	ظرفیت اصلی شبکه برای یادگیری الگوهای اولیه در داده‌های منابع آب.
۳	تعداد نرون‌های لایه دوم (Neurons in Hidden Layer 2)	۵-۱۰۰	۲۴	برای پالایش اطلاعات لایه اول و جلوگیری از پیچیدگی بیش از حد مدل.
۴	تابع فعال‌سازی (Activation Function)	{ReLU, Tanh, Sigmoid}	Tanh	تابعی که ماهیت غیرخطی به مدل می‌دهد تا بتواند نوسانات هیدرولوژیکی را شبیه‌سازی کند.
۵	نرخ یادگیری (Learning Rate)	۰/۱ - ۰/۰۰۱	۰.۰۱۲	گام‌های حرکت الگوریتم به سمت کمینه خطا؛ مقدار بهینه مانع از واگرایی مدل می‌شود.
۶	الگوریتم بهینه‌ساز وزن‌ها (Optimizer)	Adam, SGD, {RMSprop}	Adam	روشی که وزن‌های شبکه را در هر تکرار برای رسیدن به بالاترین دقت به‌روزرسانی می‌کند.
۷	تعداد تکرار/اپوک (Number of Epochs)	۱۰۰-۱۰۰۰	۵۰۰	تعداد دفعات مرور کامل مجموعه داده برای اطمینان از همگرایی مدل.
۸	اندازه دسته‌بندی (Batch Size)	۱۶-۱۲۸	۳۲	تعداد نمونه‌های داده که در هر مرحله از آموزش به طور همزمان پردازش می‌شوند.
۹	روش مقداردهی اولیه وزن‌ها (Weight Initialization)	Uniform, Normal, {Xavier}	Xavier	روشی برای جلوگیری از اشباع شدن نرون‌ها در مراحل اولیه آموزش.
۱۰	نسبت داده‌های اعتبارسنجی (Validation Split)	۰/۱ - ۰/۳	۰.۲۰	درصدی از داده‌ها که برای کنترل لحظه‌ای خطای مدل و جلوگیری از بیش‌برازش استفاده شد.
مدل جنگل تصادفی (RF)				
۱	تعداد درخت‌های تصمیم (n_estimators)	۵۰-۵۰۰	۲۰۰	تعداد کل درخت‌ها در جنگل؛ افزایش آن باعث پایداری پیش‌بینی و کاهش واریانس می‌شود.
۲	حداکثر عمق درخت (Max Depth)	۵-۵۰	۱۵	محدود کردن رشد درخت برای جلوگیری از بیش‌برازش داده‌ها و بهبود تعمیم‌دهی.

۳	حداقل نمونه برای تقسیم (Min Samples Split)	۲-۳۰	۵	حداقل تعداد داده در یک گره برای ایجاد انشعاب؛ مانع از ایجاد شاخه‌های بسیار جزئی می‌شود.
۴	حداقل نمونه در برگ (Min Samples Leaf)	۱-۱۰	۲	حداقل تعداد نمونه‌ای که باید در هر گره نهایی (برگ) باقی بماند تا تقسیم تایید شود.
۵	تعداد ویژگی‌های تصادفی (Max Features)	{Auto, Sqrt, Log2}	Sqrt	تعداد متغیرهای ورودی که در هر تقسیم به طور تصادفی بررسی می‌شوند تا درخت‌ها متفاوت باشند.
۶	روش نمونه‌گیری (Bootstrap)	{True, False}	True	استفاده از نمونه‌گیری مجدد با جایگذاری برای ایجاد مجموعه‌داده‌های متفاوت برای هر درخت.
۷	معیار سنجش ناخالصی (Criterion)	{Gini, Entropy, MSE}	MSE	شاخصی برای اندازه‌گیری کیفیت تقسیم‌ها (در اینجا بر اساس کمینه کردن میانگین مربعات خطا).
۸	حداکثر گره‌های پایانی (Max Leaf Nodes)	Unlimited / 10 - 100	Unlimited	محدود کردن تعداد کل برگ‌ها در هر درخت برای کنترل پیچیدگی کلی مدل.
یادگیری عمیق شبکه حافظه طولانی-کوتاه مدت (LSTM)				
۱	تعداد واحدهای حافظه (LSTM Units/Cells)	۱۰-۲۰۰	۸۰	تعداد سلول‌های بازگشتی در لایه پنهان که مسئول نگهداری اطلاعات در گام‌های زمانی مختلف هستند.
۲	طول گام زمانی (Look-back / Time Steps)	۱-۱۲	۳	تعداد گام‌های زمانی گذشته که مدل برای پیش‌بینی گام آینده به آن‌ها نگاه می‌کند.
۳	نرخ حذف (Dropout Rate)	۰/۵ - ۰/۱	۰/۲	درصدی از نرون‌ها که به طور تصادفی در هر تکرار حذف می‌شوند تا از وابستگی بیش از حد مدل به داده‌های آموزش جلوگیری شود.
۴	نرخ یادگیری (Learning Rate)	۰/۱ - ۰/۰۰۰۱	۰/۰۰۱	تعیین‌کننده سرعت و دقت مدل در به‌روزرسانی وزن‌ها بر اساس الگوریتم گرادیان نزولی
۵	اندازه دسته‌بندی (Batch Size)	۸-۱۲۸	۳۲	تعداد نمونه‌های داده که قبل از به‌روزرسانی پارامترهای داخلی مدل، به صورت یکجا پردازش می‌شوند.
۶	تعداد اپوک (Number of Epochs)	۵-۵۰۰	۲۰۰	تعداد دفعات کامل عبور مجموعه‌داده از شبکه برای رسیدن به وضعیت پایدار یادگیری.
۷	تابع فعال‌سازی (Activation Function)	Tanh, ReLU, {Sigmoid}	Tanh	تابعی که در داخل سلول‌های LSTM برای کنترل جریان اطلاعات و جلوگیری از انفجار گرادیان استفاده می‌شود.
۸	الگوریتم بهینه‌ساز (Optimizer)	Adam, RMSprop, {SGD}	Adam	الگوریتم محاسباتی که نرخ یادگیری را به صورت تطبیقی تنظیم می‌کند تا همگرایی سریع‌تر حاصل شود.
مدل بهینه سازی عقاب (AO)				
۱	تعداد جمعیت (Population Size - N)	۳۰-۱۰۰	۵۰	تعداد عوامل جستجو (عقاب‌ها) که فضای مسئله را برای یافتن بهترین هایپرپارامترها پیمایش می‌کنند.

۲	حداکثر تکرار (Max Iterations - T)	۱۰۰-۵۰۰	۲۰۰	حد نهایی اجرای الگوریتم برای تضمین همگرایی به بهینه جهانی.
۳	ضریب تنظیم استخراج (α - Alpha)	۰/۹ - ۰/۱	۰/۱	پارامتر کنترل کننده برای تنظیم فاز استخراج؛ مقدار ۰/۱ در اکثر توابع تست بهترین نتیجه را داده است.
۴	ضریب تعدیل استخراج (δ - Delta)	۰/۹ - ۰/۱	۰/۱	پارامتری که گام‌های جستجو را در فازهای نهایی حمله به طعمه (بهینه‌سازی) تنظیم می‌کند.
۵	ثابت نوسان (Fixed Value - U)	Fixed	۰/۰۰۵۶۵	مقدار ثابت کوچکی که در محاسبات مربوط به جستجوی مارپیچی استفاده می‌شود.
۶	محدوده چرخه جستجو (r_1 - Search Cycle)	۱-۲۰	۱۰	مقداری که برای تثبیت تعداد چرخه‌های جستجو در فازهای اکتشافی تعیین می‌گردد.

۳- نتایج و بحث

در این پژوهش، عملکرد سه مدل LSTM-AO، ANN-AO و RF-AO در دو فاز آموزش و آزمون و تحت هفت سناریوی ورودی مختلف ارزیابی شد. همان گونه که در جداول (۵) تا (۷) مشاهده می‌شود، کارایی مدل‌ها با استفاده از شاخص‌های R^2 ، MAE، RMSE و NSE بررسی شده است. نتایج نشان داد که نوع سناریوی ورودی اثر مستقیمی بر دقت شبیه‌سازی و توان تعمیم‌پذیری مدل‌ها داشته و اختلاف معنی‌داری میان مدل‌ها از نظر پایداری و صحت پیش‌بینی مشاهده می‌شود. بر اساس نتایج جدول (۶)، مدل ANN-AO در فاز آموزش عملکرد نسبتاً مناسبی داشت و در اغلب سناریوها مقادیر R^2 بیش از ۰/۸۳ به دست آمد. در این میان، سناریوی SN_4 با ساختار ۴-۲۰-۱ بهترین عملکرد را در مرحله آموزش نشان داد. با این حال، در فاز آزمون، نوسان عملکرد این مدل بیشتر بود؛ به طوری که سناریوی SN_3 با افت محسوس شاخص‌ها، بیانگر ضعف مدل در تعمیم‌پذیری تحت این ترکیب ورودی است. در مقابل، سناریوی SN_5 در مرحله آزمون با $R^2=0.879$ و $NSE=0.652$ مناسب‌ترین وضعیت را در میان سناریوهای این مدل نشان داد. نتایج جدول (۷) بیانگر آن است که مدل RF-AO در مقایسه با سایر مدل‌ها از پایداری بیشتری برخوردار بوده است. در فاز آموزش، تمامی سناریوها عملکرد مطلوبی داشتند و مقادیر R^2 در بازه ۰/۸۷۰ تا ۰/۹۰۵ قرار گرفت. همچنین در فاز آزمون نیز تغییرات شاخص‌ها محدودتر بود و سناریوهای SN_2 و SN_4 بهترین نتایج را ارائه کردند. به طور کلی، تداوم مقادیر نسبتاً بالای R^2 و NSE در کنار خطاهای کمتر در دو فاز آموزش و آزمون، نشان‌دهنده قابلیت مناسب این مدل در شبیه‌سازی رواناب و پایداری بیشتر آن در برابر تغییر سناریوهای ورودی است. مطابق جدول (۸)، مدل LSTM-AO نیز عملکرد قابل قبولی در شبیه‌سازی نشان داد، هرچند کارایی آن وابستگی محسوسی به نوع سناریوی ورودی داشت. در فاز آموزش،

سناریوی SN_3 بهترین نتیجه را ثبت کرد، در حالی که در فاز آزمون، سناریوهای SN_1 و SN_6 از نظر شاخص‌های ارزیابی مناسب‌تر بودند. در مقابل، کاهش دقت در سناریوی SN_7 نشان داد که عملکرد این مدل در همه ترکیب‌های ورودی یکنواخت نیست. در مجموع، با مقایسه نتایج جدول‌های (۶) تا (۸) می‌توان بیان کرد که مدل RF-AO از نظر ثبات و تعمیم‌پذیری، عملکرد مناسب‌تری نسبت به دو مدل دیگر داشته است، در حالی که LSTM-AO و ANN-AO در برخی سناریوهای خاص نتایج قابل قبولی ارائه کرده‌اند. در شکل (۴)، پراکنش مقادیر برآوردشده و مشاهده‌شده برای مدل ANN-AO نشان می‌دهد که مدل در بازتولید روند کلی داده‌ها موفق بوده، اما در مقادیر حدی و رخدادهای اوج دبی، پراکندگی نقاط نسبت به خط یک‌به‌یک افزایش یافته است. شیب رگرسیون کمتر از یک و وجود انحراف از خط مرجع، بیانگر تمایل مدل به کم‌برآوردی نسبی در دبی‌های بالا و افت دقت در بازنمایی رخدادهای شدید است. با این حال، مقدار بالای ضریب همبستگی گزارش‌شده در هر دو فاز آموزش و آزمون، نشان می‌دهد که این مدل توانسته الگوی کلی تغییرات رواناب را تا حد قابل قبولی شبیه‌سازی کند. در شکل (۶)، مدل RF-AO عملکرد بهتری نسبت به ANN-AO از خود نشان داده است. در هر دو فاز آموزش و آزمون، داده‌ها نسبتاً نزدیک به خط یک‌به‌یک قرار گرفته‌اند و شیب رگرسیون نیز به مقدار ایده‌آل نزدیک‌تر است. این الگو نشان می‌دهد که مدل جنگل تصادفی در شبیه‌سازی تغییرات دبی، به‌ویژه در بازه‌های میانی و بالا، از پایداری و تعمیم‌پذیری بیشتری برخوردار بوده است. همچنین، پراکندگی کمتر نقاط در فاز آزمون نسبت به مدل‌های مبتنی بر شبکه عصبی، حاکی از مقاومت بیشتر این مدل در برابر بیش‌برازش و نوسانات داده‌های هیدرولوژیکی است. مطابق شکل (۷)، سری زمانی پیش‌بینی‌شده توسط RF-AO در هر دو فاز آموزش و آزمون، تطابق مناسبی با سری زمانی مشاهداتی دارد و به‌ویژه در بازه‌های اوج، الگوی تغییرات دبی را با دقت قابل قبول دنبال می‌کند. در نمای

نسبت به ترکیب ورودی‌ها حساس‌تر بوده و احتمال بیش‌برازش در آن بیشتر است. این الگو در مدل LSTM-AO نیز مشاهده می‌شود، با این تفاوت که این مدل توانسته وابستگی‌های زمانی داده‌های هیدرولوژیکی را بهتر بازنمایی کند، اما در بازتولید برخی رخدادهای اوج، دقت آن نسبت به RF-AO پایین‌تر بوده است. از منظر هیدرولوژیکی، برتری RF-AO را می‌توان ناشی از ماهیت ensemble و توان بالای آن در مدل‌سازی روابط غیرخطی و ناهمگون سری‌های زمانی رواناب دانست. رفتار این مدل در شکل‌های پراکنش و سری زمانی نشان می‌دهد که داده‌های پیش‌بینی‌شده در بیشتر موارد به خط یک‌به‌یک نزدیک بوده و هم‌زمان در بازتولید روند کلی و قله‌های هیدروگراف نیز موفق عمل کرده است. همچنین، هیستوگرام خطا در شکل (۸) نشان می‌دهد که خطاها حول صفر متمرکزند و سوگیری سیستماتیک چشمگیری وجود ندارد؛ بنابراین، مدل نه تنها از دقت مناسب، بلکه از تعادل آماری مطلوبی نیز برخوردار است. این ویژگی برای کاربردهای منابع آب، به‌ویژه در پیش‌بینی دبی و مدیریت سیلاب، اهمیت زیادی دارد؛ زیرا مدل باید علاوه بر دقت بالا، در شرایط داده‌های دیده‌نشده نیز رفتار باثباتی داشته باشد. در مقابل، نتایج ANN-AO نشان می‌دهد که این مدل هرچند قادر به یادگیری الگوی کلی داده‌ها بوده، اما در برخی سناریوها، به‌ویژه در آزمون، افت محسوسی در عملکرد داشته است (جدول ۶). این مسئله را می‌توان به حساسیت شبکه عصبی پیش‌خور به تعداد و ماهیت ورودی‌ها، همچنین وابستگی آن به تنظیمات بهینه وزن‌ها و بایاس‌ها نسبت داد. از سوی دیگر، LSTM-AO به دلیل ساختار حافظه‌دار خود، انتظار می‌رود برای داده‌های سری زمانی هیدرولوژیکی مناسب باشد، و نتایج جدول (۸) نیز این موضوع را تا حدی تأیید می‌کند؛ با این حال، کاهش دقت در برخی سناریوها و هموارسازی نسبی پیک‌ها در شکل‌های سری زمانی نشان می‌دهد که این مدل در بازنمایی رخدادهای تند و ناگهانی هنوز محدودیت دارد. در واقع، اگرچه LSTM وابستگی‌های زمانی را بهتر می‌آموزد، اما برای داده‌هایی با نوسانات شدید و رفتار غیرایستا، لزوماً بهترین گزینه نیست. در مجموع، هم‌خوانی نتایج کمی جداول و شواهد کیفی شکل‌ها نشان می‌دهد که انتخاب سناریوی ورودی نقش تعیین‌کننده‌ای در کیفیت شبیه‌سازی دارد و صرف انتخاب یک الگوریتم پیش‌بینی، تضمین‌کننده دقت بالا نیست. از این منظر، سناریوهایی که در آن‌ها شاخص‌های دقت بالا و خطای کمتر هم‌زمان مشاهده شده‌اند، گزینه‌های مناسب‌تری برای کاربردهای عملی محسوب می‌شوند. بر همین اساس، RF-AO را می‌توان پایدارترین و قابل‌اعتمادترین مدل در این مطالعه دانست، در حالی که ANN-AO و LSTM-AO در برخی ترکیب‌های ورودی عملکرد قابل قبولی داشته‌اند اما از نظر ثبات و تعمیم‌پذیری در سطح پایین‌تری قرار گرفته‌اند. این نتیجه‌گیری با ماهیت پیچیده، غیرخطی و رخدادمحور داده‌های هیدرولوژیکی سازگار است و نشان می‌دهد که بهره‌گیری از

بزرگ‌نمایی‌شده نیز هم‌پوشانی مناسب میان داده‌های مشاهده‌ای و پیش‌بینی‌شده دیده می‌شود، هرچند در برخی رخدادهای تیز و ناگهانی، میزان هموارسازی پیش‌بینی‌ها باعث کاهش جزئی در بازنمایی قله‌ها شده است. از دیدگاه کاربردهای هیدرولوژیکی، این رفتار قابل انتظار است، زیرا مدل‌های داده‌محور معمولاً در مواجهه با پیک‌های بسیار تند، مقداری از جزئیات نوسانی را از دست می‌دهند، اما در عوض تصویر پایدارتری از روند کلی ارائه می‌کنند. در شکل (۸)، هیستوگرام خطا و برازش نرمال مدل RF-AO نشان می‌دهد که خطاها در هر دو فاز آموزش و آزمون حول صفر متمرکز هستند و سوگیری سیستماتیک چشمگیری وجود ندارد. در فاز آموزش، میانگین خطا تقریباً صفر و انحراف معیار کمتر است که دلالت بر دقت بالاتر و تمرکز بیشتر خطاها دارد. در فاز آزمون نیز اگرچه پراکندگی خطاها افزایش یافته، اما میانگین خطا همچنان نزدیک به صفر باقی مانده است. این نتیجه از منظر منابع آب اهمیت دارد، زیرا نشان می‌دهد مدل علاوه بر دقت مناسب، از نظر جهت خطا نیز رفتار متعادلی داشته و تمایل مشخصی به بیش‌برآورد یا کم‌برآورد ندارد. در شکل (۹)، پراکنش نقاط مدل LSTM-AO در فازهای آموزش و آزمون نشان می‌دهد که این مدل نیز قادر به یادگیری الگوهای غیرخطی سری زمانی بوده است، اما پراکندگی نقاط نسبت به مدل RF-AO بیشتر است. همبستگی بالای داده‌ها با خط رگرسیون، بیانگر توان مناسب LSTM در استخراج وابستگی‌های زمانی است، ولی فاصله بیشتر برخی نقاط از خط یک‌به‌یک، به‌ویژه در دبی‌های زیاد، نشان می‌دهد که در برخی رخدادهای حدی، دقت مدل کاهش یافته است. بنابراین، LSTM-AO از نظر ساختار زمانی مزیت دارد، اما در این مطالعه از نظر پایداری کلی به پای مدل RF-AO نرسیده است. در شکل (۱۰)، نمودارهای سری زمانی مدل LSTM-AO نشان می‌دهد که این مدل در بازتولید روند عمومی دبی موفق بوده و در بسیاری از بازه‌ها، تغییرات مشاهداتی را دنبال کرده است. با این حال، در برخی قله‌های تیز و کوتاه‌مدت، پاسخ مدل هموارتر از داده واقعی است و همین امر موجب کاهش دقت در ثبت اوج‌ها شده است. این ویژگی در مدل‌های حافظه‌دار رایج است، زیرا LSTM معمولاً الگوهای وابستگی زمانی را به‌خوبی یاد می‌گیرد، اما در برابر نوسانات ناگهانی و بسیار محلی، واکنشی محافظه‌کارانه‌تر نشان می‌دهد. یافته‌های این مطالعه نشان می‌دهد که عملکرد مدل‌ها به‌طور معنی‌داری تحت تأثیر نوع ساختار ورودی و معماری مدل قرار گرفته است. به‌طور کلی، مقایسه شاخص‌های R^2 ، MAE، RMSE و NSE در جداول (۶) تا (۸) بیانگر آن است که مدل RF-AO نسبت به دو مدل دیگر از پایداری و تعمیم‌پذیری بیشتری برخوردار بوده است. این موضوع به‌ویژه در فاز آزمون آشکار است، جایی که افت عملکرد این مدل نسبت به آموزش محدود بوده و مقادیر خطا در سطح قابل قبولی باقی مانده‌اند. در مقابل، مدل ANN-AO اگرچه در برخی سناریوها عملکرد آموزشی مناسبی داشته، اما نوسان بیشتر در مرحله آزمون نشان می‌دهد که این مدل

مدل های ensemble، به ویژه در ترکیب با بهینه سازی مناسب، می تواند راهبرد مؤثرتری برای پیش بینی متغیرهای منابع آب باشد.

۴- نتیجه گیری نهایی

این پژوهش با هدف توسعه یک سیستم پیش بینی روزانه جریان رودخانه در حوضه آبریز سیمینه رود، عملکرد سه الگوریتم داده محور شامل شبکه عصبی مصنوعی پرسپترون چندلایه، جنگل تصادفی و حافظه طولانی کوتاه مدت را که همگی با الگوریتم فراابتکاری عقاب بهینه سازی شده بودند، تحت هفت سناریوی مختلف ورودی مورد ارزیابی قرار داد. نتایج به دست آمده نشان داد که الگوریتم بهینه سازی عقاب کارایی مناسبی در تنظیم هایپرپارامترهای مدل ها داشته است، اما توانایی نهایی مدل ها در تعمیم پذیری به طور مستقیم به ساختار ریاضی آن ها و گزینش سناریوی ورودی بستگی دارد. در این میان، مدل جنگل تصادفی بهینه شده با الگوریتم عقاب در سناریوی چهارم با ترکیب متغیرهای هیدروکلیماتیک و تأخیرهای زمانی دبی، برترین عملکرد را در فاز آزمون ارائه کرد و به ضریب تبیین معادل ۰/۸۳۴، ضریب کارایی ناش-ساتکلیف معادل ۰/۸۸ و ریشه میانگین مربع خطا معادل ۱۱/۹۳۱ مترمکعب بر ثانیه دست یافت که نشان دهنده بهبود

کارایی ناش-ساتکلیف آن نسبت به مدل های حافظه طولانی کوتاه مدت و شبکه عصبی مصنوعی به ترتیب به میزان ۲/۳ درصد و ۳۶ درصد است. تحلیل آماری و بصری توزیع خطاها و نمودارهای پراکنش تأیید کرد که مدل جنگل تصادفی به دلیل ساختار مبتنی بر یادگیری جمعی، پایداری بالاتری در مواجهه با نوسانات دبی روزانه و پیشگیری از بیش برآزش دارد. در مقابل، مدل شبکه عصبی مصنوعی با وجود سرعت بالا، حساسیت شدیدی به تغییر سناریوهای ورودی نشان داد و مدل یادگیری عمیق حافظه طولانی کوتاه مدت نیز علی رغم بازنمایی مناسب وابستگی های زمانی، در شبیه سازی پیک های سیلابی ناگهانی رفتار محافظه کارانه ای داشت و منجر به هموارسازی دبی های حدی شد. دستاورد کاربردی این تحقیق نشان می دهد که افزایش بی رویه متغیرهای ورودی لزوماً به بهبود پیش بینی منجر نمی شود و تعیین دقیق ترکیب بهینه ورودی ها با استفاده از شاخص های همبستگی نقشی کلیدی در کاهش عدم قطعیت های مدل سازی دارد. در نهایت، مدل توسعه یافته جنگل تصادفی پیوند یافته با الگوریتم عقاب به عنوان ابزاری پایدار و دقیق برای پیش بینی جریان روزانه پیشنهاد می شود که می تواند در سامانه های هوشمند تصمیم یار مدیریت منابع آب و پیش بینی سیلاب های حوضه سیمینه رود مورد استفاده قرار گیرد.

جدول ۶- نتایج عملکرد مدل ANN-AO در دو فاز آموزش و آزمون

Table 6. Performance results of the ANN-AO model in the training and testing phases.

سناریوها				ساختار شبکه				فرایند آموزش (۷۰ درصد)				فرایند آزمون (۳۰ درصد)			
SN ₁	SN ₂	SN ₃	SN ₄	SN ₅	SN ₆	SN ₇		R ²	MAE (m ³ /s)	RMSE(m ³ /s)	NSE (%)	R ²	MAE (m ³ /s)	RMSE (m ³ /s)	NSE (%)
۱-۲۰-۱	۱-۲۰-۲	۱-۲۰-۳	۱-۲۰-۴	۱-۲۰-۵	۱-۲۰-۶	۱-۲۰-۷		۰.۸۶۶	۰.۵۹۰	۱.۲۳۸	۰.۶۳۴	۰.۸۵۸	۰.۵۷۹	۱.۲۶۹	۰.۶۲۴
								۰.۸۷۶	۰.۵۷۳	۱.۲۰۰	۰.۶۴۸	۰.۸۵۳	۰.۵۸۳	۱.۲۵۷	۰.۶۱۶
								۰.۸۳۳	۰.۶۱۶	۱.۳۸۱	۰.۵۹۱	۰.۴۴۹	۱.۸۸۷	۲.۵۰۹	۰.۲۵۸
								۰.۸۸۱	۰.۵۷۶	۱.۱۶۴	۰.۶۵۵	۰.۸۵۶	۰.۵۹۸	۱.۳۰۱	۰.۶۲۰
								۰.۸۴۳	۰.۵۹۸	۱.۳۳۸	۰.۶۰۵	۰.۸۷۹	۰.۵۷۸	۱.۱۸۶	۰.۶۵۲
								۰.۸۷۶	۰.۵۹۰	۱.۱۹۲	۰.۶۴۸	۰.۸۷۵	۰.۵۸۵	۱.۱۸۵	۰.۶۴۷
								۰.۸۸۰	۰.۶۰۸	۱.۱۷۲	۰.۶۵۴	۰.۸۵۲	۰.۶۵۳	۱.۲۹۸	۰.۶۱۵

جدول ۷- نتایج عملکرد مدل RF-AO در دو فاز آموزش و آزمون

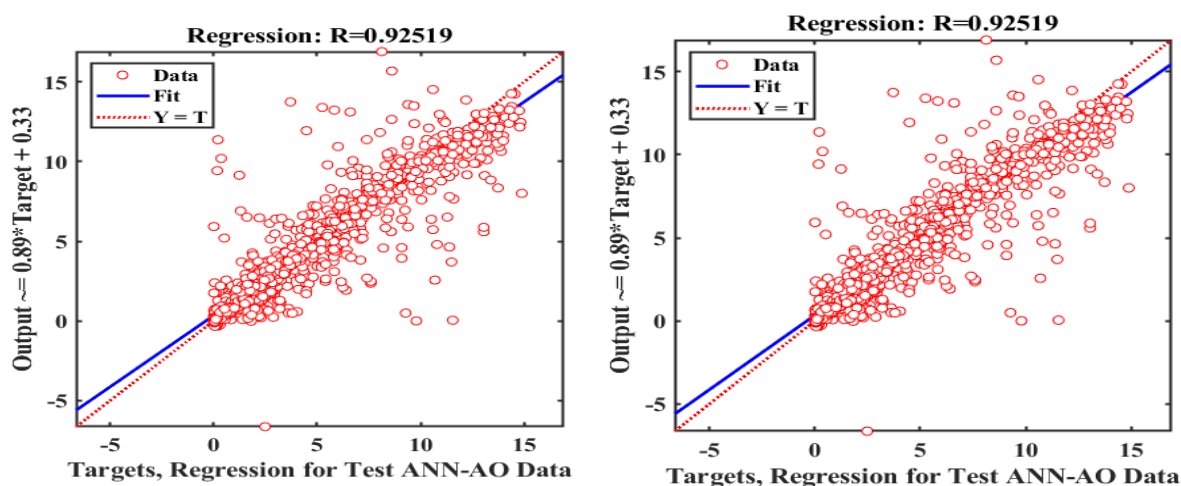
Table 7. Performance results of the RF-AO model in the training and testing phases.

سناریوها	فرایند آموزش (۷۰ درصد)				فرایند آزمون (۳۰ درصد)			
	R ²	MAE (m ³ /s)	RMSE(m ³ /s)	NSE (%)	R ²	MAE (m ³ /s)	RMSE (m ³ /s)	NSE (%)
SN ₁	۰.۸۷۰	۳.۰۲۴	۱۰.۳۰۷	۰.۹۰۰	۰.۸۲۴	۳.۵۶۹	۱۲.۳۱۸	۰.۸۸۸
SN ₂	۰.۸۹۲	۲.۵۸۶	۹.۴۱۶	۰.۹۰۲	۰.۸۳۴	۳.۲۵۶	۱۱.۹۳۴	۰.۸۸۲
SN ₃	۰.۸۸۹	۲.۶۳۱	۹.۶۱۳	۰.۸۸۴	۰.۸۲۸	۳.۳۳۳	۱۲.۱۴۵	۰.۸۶۶
SN ₄	۰.۹۰۱	۲.۴۴۴	۹.۰۳۶	۰.۹۰۴	۰.۸۳۴	۳.۲۴۹	۱۱.۹۳۱	۰.۸۸۰
SN ₅	۰.۸۹۸	۲.۴۹۸	۹.۲۰۵	۰.۸۹۴	۰.۸۳۲	۳.۳۱۵	۱۲.۰۱۹	۰.۸۷۰
SN ₆	۰.۸۹۵	۲.۵۵۲	۹.۳۵۰	۰.۸۸۶	۰.۸۲۸	۳.۳۹۳	۱۲.۱۴۵	۰.۸۶۲
SN ₇	۰.۹۰۵	۲.۴۰۹	۸.۹۰۳۱	۰.۸۹۳	۰.۸۲۹	۳.۴۰۷	۱۲.۱۳۸	۰.۸۶۵

جدول ۸- نتایج عملکرد مدل LSTM-AO در دو فاز آموزش و آزمون

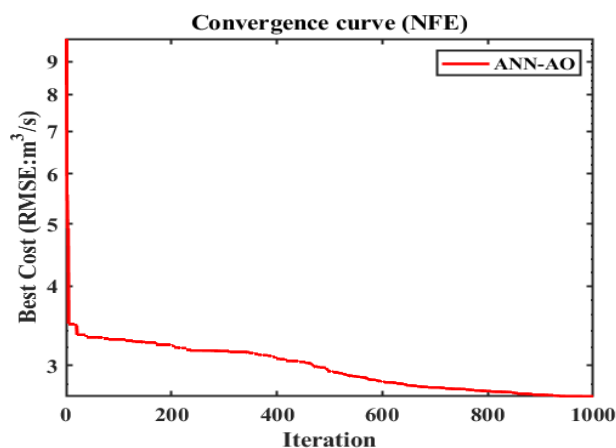
Table 8. Performance results of the LSTM-AO model in the training and testing phases.

فرایند آزمون (۳۰ درصد)				فرایند آموزش (۷۰ درصد)				سناریوها
NSE (%)	RMSE (m ³ /s)	MAE (m ³ /s)	R ²	NSE (%)	RMSE(m ³ /s)	MAE (m ³ /s)	R ²	
۰.۸۶۰	۱۰.۵۴۱	۳.۵۵۸	۰.۸۶۲	۰.۸۳۶	۱۱.۷۸۲	۳.۵۸۷	۰.۸۳۷	SN ₁
۰.۷۹۷	۱۲.۱۵۷	۴.۲۱۹	۰.۸۱۸	۰.۸۱۴	۱۲.۷۶	۴.۱۱۱	۰.۸۳۹	SN ₂
۰.۸۳۶	۱۱.۲۷۵	۳.۷۶۴	۰.۸۳۷	۰.۸۵۸	۱۱.۰۱۳	۳.۵۵۰	۰.۸۵۸	SN ₃
۰.۸۴۳	۱۱.۴۵۳	۴.۲۱۳	۰.۸۴۴	۰.۸۵۲	۱۱.۰۸۶	۳.۹۴۰	۰.۸۵۳	SN ₄
۰.۸۴۴	۱۱.۱۹۱	۴.۰۱۹	۰.۸۴۵	۰.۸۴۴	۱۱.۴۸۲	۳.۸۵۴	۰.۸۴۸	SN ₅
۰.۸۵۴	۱۰.۹۱	۴.۳۱۱	۰.۸۶۵	۰.۸۱۰	۱۲.۶۳۷	۴.۳۴۲	۰.۸۲۰	SN ₆
۰.۷۹۳	۱۲.۲۸۲	۵.۸۴۴	۰.۷۹۸	۰.۸۲۲	۱۲.۱۲۶	۵.۴۸۱	۰.۸۲۷	SN ₇



شکل ۴- نمودار پراکندگی مدل ANN-AO در دو فاز آموزش و آزمون

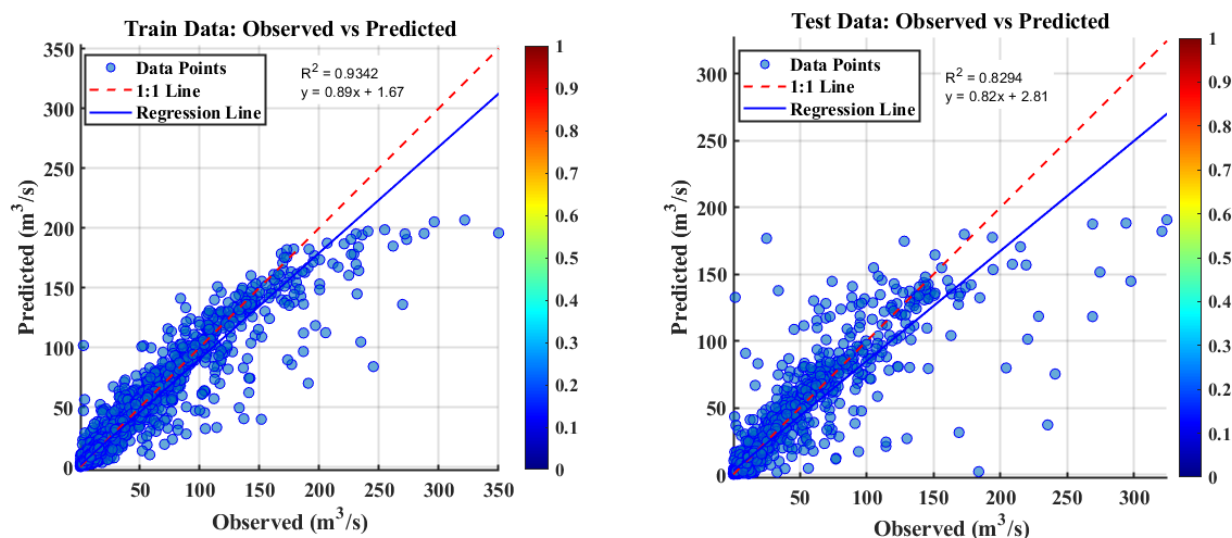
Fig 4. Scatter plots of the ANN-AO model in the training and testing phases



شکل ۵- نمودار همگرایی (NFE) مدل هیبریدی ANN-AO

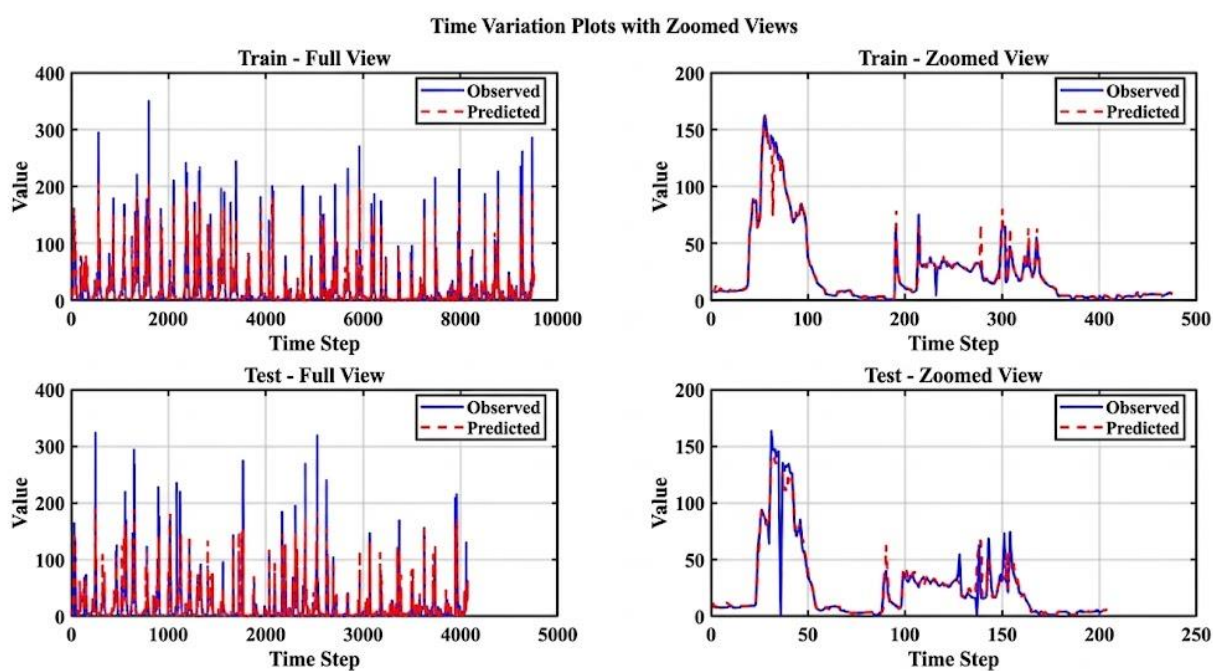
Fig 5. Convergence plot (NFE) of the hybrid ANN-AO model

Scatter Plots (RF-AO)



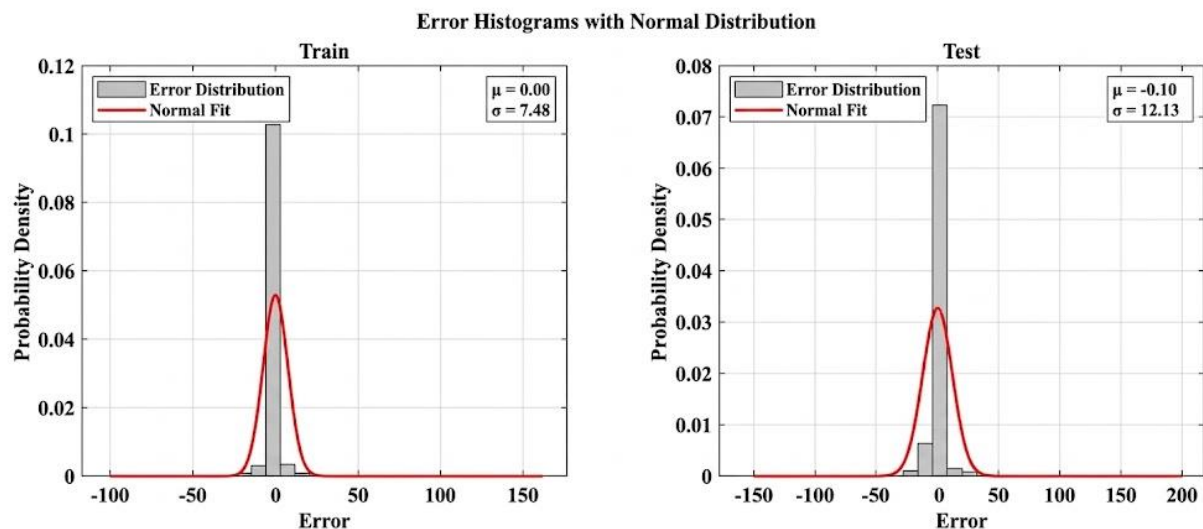
شکل ۶- نمودار پراکندگی مدل RF-AO در دو فاز آموزش و آزمون

Fig 6. Scatter plots of the RF-AO model in the training and testing phases.



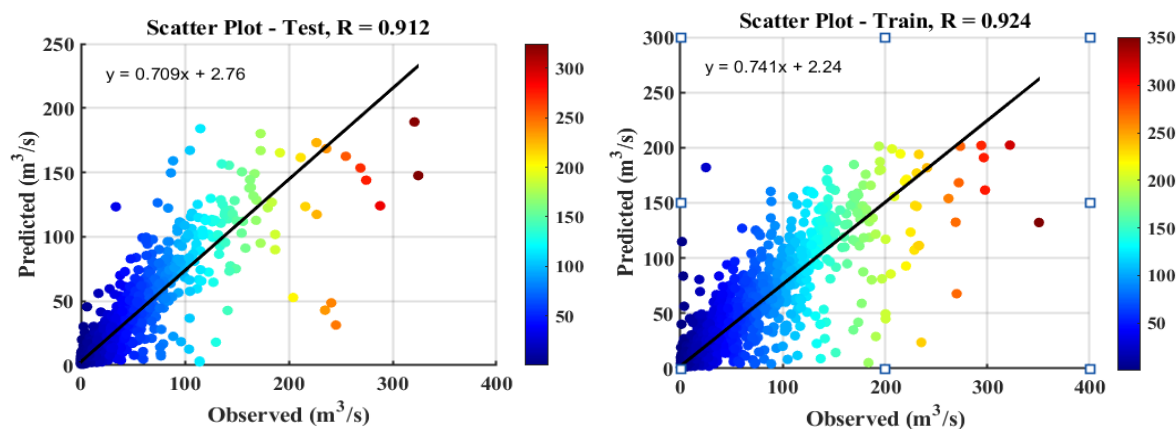
شکل ۷- نمودار سری زمانی مدل RF-AO در دو فاز آموزش و آزمون

Fig 7. Time series plots of the RF-AO model in the training and testing phase



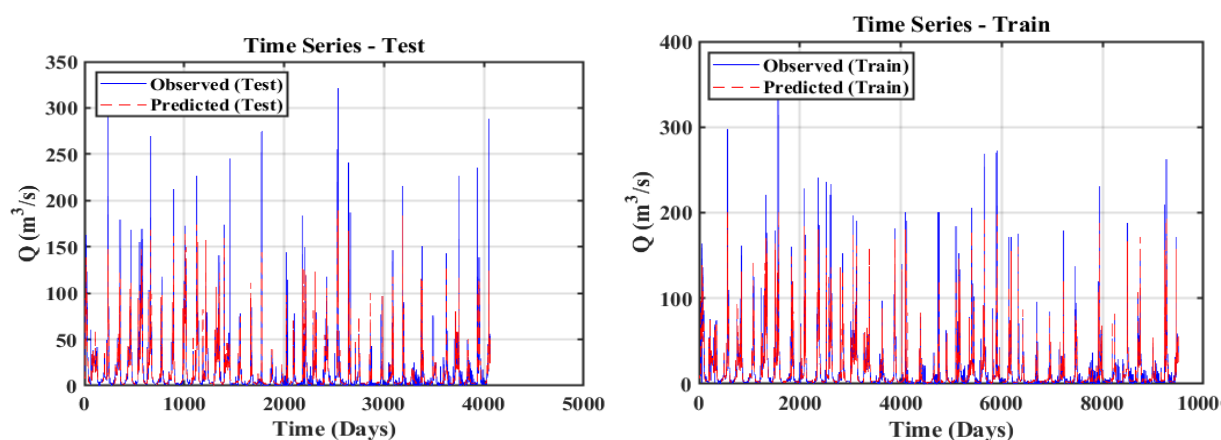
شکل ۸- منحنی توزیع نرمال و هیستوگرام خطا مدل RF-AO در دو فاز آموزش و آزمون

Fig 8. Normal probability plot and error histogram of the RF-AO model in the training and testing phases



شکل ۹- نمودار پراکندگی مدل LSTM-AO در دو فاز آموزش و آزمون

Fig 9. Scatter plots of the LSTM-AO model in the training and testing phases



شکل ۱۰- نمودار سری زمانی مدل LSTM-AO در دو فاز آموزش و آزمون

Fig 10. Time series plots of the LSTM-AO model in the training and testing phases

منابع

- (2021). Aquila of machine learning and deep learning models for the daily river streamflow forecasting. *Water Resources Management*, 39(4), 1911–1930. <https://doi.org/10.1007/s11269-024-04052-y>
- Guo, J., Liu, Y., Zou, Q., Ye, L., Zhu, S., & Zhang, H. (2023). Study on optimization and combination strategy of multiple daily runoff prediction models coupled with physical mechanism and LSTM. *Journal of Hydrology*, 624, 129969. <https://doi.org/10.1016/j.jhydrol.2023.129969>
- Hochreiter, S., & Schmidhuber, J. (1997). Long short-term memory. *Neural Computation*, 9(8), 1735–1780. <https://doi.org/10.1162/neco.1997.9.8.1735>
- Huang, J., Chen, J., Huang, H., & Cai, X. (2025). Deep Learning-Based Daily Streamflow Prediction Model for the Hanjiang River Basin. *Hydrology*, 12(7), 168. <https://doi.org/10.3390/hydrology12070168>
- Janiesch, C., Zschech, P., & Heinrich, K. (2021). Machine learning and deep learning. *Electronic Markets*, 31(3), 685–695.
- Jimmy, S. R., Sahoo, A., Samantaray, S., & Ghose, D. K. (2020). Prophecy of runoff in a river basin using various neural networks. In *Communication Software and Networks: Proceedings of INDIA 2019* (pp. 709–718). Springer. https://doi.org/10.1007/978-981-15-5397-4_72
- Le, X.-H., Ho, H. V., Lee, G., & Jung, S. (2019). Application of long short-term memory (LSTM) neural network for flood forecasting. *Water*, 11(7), 1387. <https://doi.org/10.3390/w11071387>
- Lin, Y., Wang, D., Meng, Y., Sun, W., Qiu, J., Shangguan, W., Cai, J., Kim, Y., & Dai, Y. (2023). Bias learning improves data driven models for streamflow prediction. *Journal of Hydrology: Regional Studies*, 50, 101557. <https://doi.org/10.1016/j.ejrh.2023.101557>
- Mangukiya, N. K., & Sharma, A. (2025). Deep learning-based approach for enhancing streamflow prediction in watersheds with aggregated and intermittent observations. *Water Resources Research*, 61(1),
- Abualigah, L., Yousri, D., Abd Elaziz, M., Ewees, A. A., Al-Qaness, M. A. A., & Gandomi, A. H. optimizer: a novel meta-heuristic optimization algorithm. *Computers & Industrial Engineering*, 157, 107250. <https://doi.org/10.1016/j.cie.2021.107250>
- Ahmadaali, J., Barani, G. A., Qaderi, K., & Hessari, B. (2017). Calibration and validation of model WEAP21 for Zarrineh Rud and Simineh Rud basins. *Iranian Journal of Soil and Water Research*, 48(4), 823–839. <https://doi.org/10.22059/ijswr.2017.216989.667543>
- Akiner, M. E., Kartal, V., Guzeler, A. C., & Karakoyun, E. (2024). Exploring the applicability of the experiment-based ANN and LSTM models for streamflow estimation. *Earth Science Informatics*, 17(4), 3111–3135. <https://doi.org/10.1609/aaai.v32i1.12255>
- Bargam, B., Boudhar, A., Kinnard, C., Bouamri, H., Nifa, K., & Chehbouni, A. (2024). Evaluation of the support vector regression (SVR) and the random forest (RF) models accuracy for streamflow prediction under a data-scarce basin in Morocco. *Discover Applied Sciences*, 6(6), 306. <https://doi.org/10.1007/s42452-024-05994-z>
- Bostanmaneshrad, F., Partani, S., Noori, R., Nachtnebel, H.-P., Berndtsson, R., & Adamowski, J. F. (2018). Relationship between water quality and macro-scale parameters (land use, erosion, geology, and population density) in the Siminehrood River Basin. *Science of the Total Environment*, 639, 1588–1600. <https://doi.org/10.1016/j.scitotenv.2018.05.244>
- Breiman, L. (2001). Random forests. *Machine Learning*, 45, 5–32. <https://doi.org/10.1023/A:1010933404324>
- Chen, L., Wang, Z., Jiang, Z., & Lin, X. (2024). Deep learning models for multi-step prediction of water levels incorporating meteorological variables and historical data. *Stochastic Environmental Research and Risk Assessment*, 1–23. <https://doi.org/10.1007/s00477-024-02766-4>
- Danesh, M., Gharehbaghi, A., Mehdizadeh, S., & Danesh, A. (2025). A comparative assessment

- perspective: the combined use of analytical tools* (Vol. 10). Springer Science & Business Media.
- Xu, W., Chen, J., Corzo, G., Xu, C., Zhang, X. J., Xiong, L., Liu, D., & Xia, J. (2024). Coupling deep learning and physically based hydrological models for monthly streamflow predictions. *Water Resources Research*, 60(2), e2023WR035618. <https://doi.org/10.1029/2023WR035618>
- Yao, Z., Wang, Z., Wang, D., Wu, J., & Chen, L. (2023). An ensemble CNN-LSTM and GRU adaptive weighting model based improved sparrow search algorithm for predicting runoff using historical meteorological and runoff data as input. *Journal of Hydrology*, 625, 129977. <https://doi.org/10.1016/j.jhydrol.2023.129977>
- Zhang, J., & Yan, H. (2023). A long short-term components neural network model with data augmentation for daily runoff forecasting. *Journal of Hydrology*, 617, 128853. <https://doi.org/10.1016/j.jhydrol.2022.128853>
- e2024WR037331. <https://doi.org/10.1029/2024WR037331>
- Ng, K. W., Huang, Y. F., Koo, C. H., Chong, K. L., El-Shafie, A., & Ahmed, A. N. (2023). A review of hybrid deep learning applications for streamflow forecasting. *Journal of Hydrology*, 625, 130141. <https://doi.org/10.1016/j.jhydrol.2023.130141>
- Qiao, X., Peng, T., Sun, N., Zhang, C., Liu, Q., Zhang, Y., Wang, Y., & Nazir, M. S. (2023). Metaheuristic evolutionary deep learning model based on temporal convolutional network, improved aquila optimizer and random forest for rainfall-runoff simulation and multi-step runoff prediction. *Expert Systems with Applications*, 229, 120616. <https://doi.org/10.1016/j.eswa.2023.120616>
- Rezaei, A., Babazadeh, H., Khosrojerdi, A., & Sarai-Tabrizi, M. (2025). Removal of sulfate pollutant from different samples of a river water using nanozeolite technology, case study: Gamasiab River, Iran. *PloS One*, 20(2), e0314480. <https://doi.org/10.1371/journal.pone.0314480>
- Sasmal, B., Hussien, A. G., Das, A., & Dhal, K. G. (2023). A comprehensive survey on aquila optimizer. *Archives of Computational Methods in Engineering*, 30(7), 4449–4476. <https://doi.org/10.1007/s11831-023-09945-6>
- Shi, W., Wang, M., Li, D., Li, X., & Sun, M. (2023). An improved method that incorporates the estimated runoff for peak discharge prediction on the Chinese Loess Plateau. *International Soil and Water Conservation Research*, 11(2), 290–300. <https://doi.org/10.1016/j.iswcr.2022.09.001>
- Vatanchi, S. M., Etemadfar, H., Maghrebi, M. F., & Shad, R. (2023). A comparative study on forecasting of long-term daily streamflow using ANN, ANFIS, BiLSTM and CNN-GRU-LSTM. *Water Resources Management*, 37(12), 4769–4785. <https://doi.org/10.21203/rs.3.rs-1443377/v1>
- Wang, Z., Xu, N., Bao, X., Wu, J., & Cui, X. (2024). Spatio-temporal deep learning model for accurate streamflow prediction with multi-source data fusion. *Environmental Modelling & Software*, 178, 106091. <https://doi.org/10.1016/j.envsoft.2024.106091>
- Wrisberg, N., de Haes, H. A. U., & Triebswetter, U. (2002). *Analytical tools for environmental design and management in a systems*