



Semnan University

Climate and Ecosystem of Arid and Semi-arid Regions

<https://ceasr.semnan.ac.ir>



Research Article

Performance Evaluation of Soft Computing-Based Methods for Daily Streamflow Prediction

Edris Maroofinia ^{a,*}, Seyyed Mehdi Esmat Sa'atloo ^a, Sedigheh Shakofti ^b and Roya Yazdanimehr ^c

^a PhD Graduate, Department of Civil Engineering, Faculty of Civil Engineering, Science and Research Branch, Islamic Azad University, Tehran, Iran

^b PhD Student, Department of Natural Resources Engineering, Faculty of Agriculture and Natural Resources, University of Hormozgan, Hormozgan, Iran

^c MSc Student, Department of Nature, Faculty of Natural Resources and Environment, Malayer University, Iran

ARTICLE INFO

Article type:

Research Full Paper

Article history:

Received: 18 september 2024

Revised: 26 november 2024

Accepted: 19 march 2025

Keywords:

Streamflow prediction,
Soft computing,
Deep learning,
Pearson correlation
coefficient,
Tajan Basin.

EXTENDED ABSTRACT

Background and Objectives: Accurate daily river flow forecasting particularly in humid catchments characterized by dynamic hydrological regimes plays a vital role in sustainable water resources management, hydrological planning, and flood risk mitigation. In recent years, soft computing-based models, including machine learning and deep learning approaches, have garnered significant attention in water engineering due to their inherent capacity to capture complex, nonlinear, nonstationary, and multivariate relationships. However, empirical evidence indicates that a systematic and comparative evaluation of advanced artificial intelligence models especially under humid catchment conditions and at daily temporal resolution remains insufficient. This research gap not only limits a comprehensive understanding of the strengths and limitations of each method but also poses serious challenges to the transferability, robustness, and operational reliability of these models in practical decision-making. To address this gap, the present study evaluates and compares the performance of machine learning and deep learning models for daily river flow prediction, with a specific focus on a humid basin subject to climate variability, thereby clarifying methodological innovations and operational applicability of these AI-driven approaches.

Materials and Methods: In this study, daily meteorological and hydrometric data spanning the period from 1969 to 2018 were first collected. Precipitation (P_t), evaporation (E_t), and river flow discharge with one- to three-day lags (Q_{t-1} to Q_{t-3}) were used as input variables. Pearson correlation coefficients were calculated between input and output variables, based on which five input scenarios and model configurations were selected. Data preprocessing (including homogenization, outlier removal, reconstruction of missing values, and normalization) was performed, and the dataset was split into training and testing subsets in a 70:30 ratio. For modeling, one machine learning algorithm Support Vector Regression (SVR) and two deep learning models Convolutional Neural Network (CNN) and Generative Adversarial Network (GAN) were employed. Model performance was

* Corresponding author: edris.marufynia1389@gmail.com

evaluated using multiple metrics: Coefficient of determination (R^2), Root Mean Square Error (RMSE), Percent Bias (PBIAS), and Kling–Gupta Efficiency (KGE). The results were compared in tabular form, and final outcomes were visually presented using scatter, time series, violin plot, and Taylor diagrams.

Results: The results indicated that SVR model achieved the best performance in Scenario 5 (SN5), yielding R^2 of 0.85, RMSE of 5.675 m^3/s , PBIAS close to zero (0.475%), and KGE of 0.877 during the testing phase. In this scenario, SVR outperformed the CNN and GAN models by 3.8% and 3.4% in terms of R^2 , and by 14.1% and 14.0% in terms of KGE, respectively. Moreover, the CNN and GAN models demonstrated acceptable performance only in the more complex scenarios (SN₃ to SN₅), while exhibiting poor results in the simpler scenarios (SN₁ and SN₂), with R^2 values below 0.04. Overall, SVR with the optimal input combination (SN₅) provided the most stable and accurate predictions, showing statistically significant superiority over the other two models across all effective scenarios.

Conclusion: In daily river discharge forecasting within highly variable basins, model input structure exerts a far greater influence than algorithmic complexity so much that excluding lagged discharge values led to complete predictive failure ($R^2 < 0.04$), even when employing deep learning architectures. Among the evaluated models, Support Vector Regression (SVR) under the optimal scenario (SN5) achieved the most accurate and unbiased performance in the testing phase, with $R^2 = 0.850$, RMSE = 5.675 m^3/s , and PBIAS = 0.475%, significantly outperforming both CNN and GAN. These findings substantiate a fundamental principle in hydroinformatics: Under conditions of limited data and high climatic stress, model selection should be guided by alignment with the hydrological character of the data, rather than by algorithmic novelty alone, a conclusion with direct implications for designing operational forecasting systems and adaptive water governance in vulnerable regions.

Cite this article: Maroofinia, E., Esmat Sa'atloo, S.M., Shakoft, S., & Yazdanimehr, R. (2025). Performance evaluation of soft computing-based methods for daily streamflow prediction. *Climate and Ecosystem of Arid and Semi-arid Regions*, 2(2), 206-226.

مقاله پژوهشی

ارزیابی عملکرد روش‌های مبتنی بر محاسبات نرم برای پیش‌بینی روزانه جریان رودخانه با تفکیک زمانی روزانه

ادریس معروفی نیا^{۱*}، سید مهدی عصمت ساعتلو^۱، صدیقه شکفتی^۲ و رویا یزدانی مهر^۳

۱- فارغ التحصیل دکتری، گروه عمران، دانشکده مهندسی عمران، دانشگاه علوم و تحقیقات تهران، تهران، ایران

۲- دانشجوی دکتری، گروه مهندسی منابع طبیعی، دانشکده کشاورزی و منابع طبیعی، دانشگاه هرمزگان، هرمزگان، ایران

۳- دانشجوی کارشناسی ارشد، گروه طبیعت، دانشکده منابع طبیعی و محیط زیست، دانشگاه ملایر، ملایر، ایران

*: نویسنده مسئول، edris.marufynia1389@gmail.com

اطلاعات مقاله	چکیده
نوع مقاله: مقاله کامل علمی- پژوهشی تاریخ دریافت: ۱۴۰۳/۰۶/۲۸ تاریخ ویرایش: ۱۴۰۳/۰۹/۰۶ تاریخ پذیرش: ۱۴۰۳/۱۲/۲۹ واژه‌های کلیدی: پیش‌بینی رودخانه، محاسبات نرم، یادگیری عمیق، ضریب پیرسون، حوضه تجن.	سابقه و هدف: پیش‌بینی دقیق جریان رودخانه در مقیاس روزانه، به‌ویژه در حوضه‌های آبریز مرطوب با رژیم هیدرولوژیک پویا، نقشی حیاتی در مدیریت پایدار منابع آب، برنامه‌ریزی‌های هیدرولوژیک و کاهش خطر سیلاب ایفا می‌کند. در سال‌های اخیر، مدل‌های مبتنی بر محاسبات نرم (یادگیری ماشین و عمیق) به دلیل توانایی ذاتی در مدل‌سازی روابط غیرخطی، ناپیدا و چندمتغیره، مورد توجه گسترده در مهندسی آب قرار گرفته‌اند. با این حال، شواهد تجربی نشان می‌دهد که ارزیابی سیستماتیک و مقایسه‌ای عملکرد مدل‌های پیشرفته هوش مصنوعی، به‌ویژه در شرایط حوضه‌های مرطوب و با دقت زمانی روزانه، همچنان ناکافی است. این شکاف تحقیقاتی، نه تنها درک جامع از مزایا و محدودیت‌های هر روش را کاهش می‌دهد، بلکه قابلیت انتقال، استحکام و اعتمادپذیری مدل‌ها را در تصمیم‌گیری‌های عملیاتی با چالش جدی مواجه می‌کند. این مطالعه، با هدف پر کردن این شکاف، عملکرد روش‌های یادگیری ماشین و عمیق را در پیش‌بینی جریان روزانه رودخانه با تمرکز بر یک حوضه مرطوب، تحت تأثیر تغییرات اقلیمی، ارزیابی و مقایسه می‌کند تا نوآوری‌های روش‌شناختی و امکان استفاده عملیاتی از این مدل‌ها را به‌وضوح مشخص نماید. مواد و روش‌ها: در این پژوهش، ابتدا داده‌های هواشناسی و ایستگاه‌های هیدرومتری در مقیاس روزانه در دوره زمانی ۱۳۴۸ لغایت پایان ۱۳۹۷ رودخانه تجن واقع در استان مازندران اخذ گردید. از داده‌های بارش (P_t)، تبخیر (E_t) دبی جریان رودخانه از یک تا سه روز تأخیر (Q_{t-1} تا Q_{t-3}) به عنوان متغیرهای ورودی استفاده گردید. بر اساس ضریب پیرسون، همبستگی بین متغیرهای ورودی و خروجی اندازه‌گیری شد و پنج سناریو و ترکیب مدل انتخاب گردید. عملیات پیش‌پردازش داده (از قبیل همگن‌سازی، حذف داده‌های پرت، بازسازی اطلاعات و داده مفقوده و نرمال‌سازی) و تفکیک داده‌ها به آموزش و آزمون با نسب ۷۰ به ۳۰ درصد روی داده‌ها انجام گرفت. جهت انجام مدل‌سازی از یک مدل یادگیری ماشین بردار پشتیبان رگرسیون (SVR) و دو مدل یادگیری عمیق شبکه عصبی پیچشی یا فراگشتی (CNN) و شبکه مولد رقابتی یا تخصی (GAN) استفاده گردید. جهت بررسی و ارزیابی میزان دقت مدل، از معیارهای ارزیابی ضریب تبیین (تعیین) (R^2)، ریشه میانگین مربعات خطا (RMSE)، درصد سوگیری یا ناریبی (PBIAS) و ضریب بهره‌وری کلینگ-گوپتا (KGE) استفاده گردید. خروجی نتایج در جداول مربوطه مورد مقایسه قرار گرفته و نتایج نهایی به صورت بصری با نمودارهای پراکندگی، سری زمانی، ویولن پلات و دیاگرام تیلور نمایش داده

شد.

یافته‌ها: نتایج نشان داد که مدل SVR در سناریوی پنجم (SN_5) بهترین عملکرد را داشته است، به‌طوری که در فاز آزمون، ضریب تبیین (R^2) برابر با ۰/۸۵، ریشه میانگین مربعات خطا (RMSE) معادل ۵/۶۷۵ متر مکعب بر ثانیه، درصد سوگیری یا بایاس (PBIAS) نزدیک به صفر (۰/۴۷۵ درصد) و کارایی کلی کمینه (KGE) برابر با ۰/۸۷۷ را به‌دست آورد. این مدل در مقایسه با CNN و GAN در همان سناریو، به‌ترتیب ۳/۸ و ۳/۴ درصد از نظر R^2 و ۱۴/۱ و ۱۴ درصد از نظر KGE در فاز آزمون برتری داشت. همچنین، مدل‌های CNN و GAN تنها در سناریوهای پیچیده (SN_3 تا SN_5) توانایی مطلوبی از خود نشان دادند، در حالی که در سناریوهای ساده (SN_1 و SN_2) عملکرد ضعیفی داشتند ($R^2 < 0.04$). در مجموع، SVR با ترکیب بهینه‌ی ورودی‌ها (SN_5) پایدارترین و دقیق‌ترین پیش‌بینی‌ها را ارائه کرد و از نظر معیارهای آماری در تمام سناریوهای مؤثر، بر دو مدل دیگر برتری معنی‌داری داشت.

نتیجه‌گیری: این مطالعه نشان داد که در پیش‌بینی روزانه‌ی دبی رودخانه در حوضه‌های با نوسانات شدید، ساختار ورودی مدل تأثیری چندمراتبی‌تر از پیچیدگی الگوریتمی دارد؛ چنان‌که حذف دبی‌های تأخیری منجر به شکست کامل پیش‌بینی ($R^2 < 0.04$) شد، حتی با استفاده از معماری‌های یادگیری عمیق. در میان مدل‌های مورد ارزیابی، ماشین بردار پشتیبان (SVR) در سناریوی بهینه (SN_5) با دستیابی به $R^2 = 0.850$ ، $RMSE = 5.675 \text{ m}^3/\text{s}$ و $PBIAS = 0.475\%$ در فاز آزمون، دقیق‌ترین و بی‌سوگیرترین عملکرد را از خود نشان داد و بر CNN و GAN برتری معنی‌داری یافت. یافته‌ها گواهی بر این اصل هیدروانفورماتیک هستند که در شرایط داده‌های محدود و پرتنش اقلیمی، انتخاب مدل باید بر اساس تطبیق با ماهیت هیدرولوژیک داده باشد، نه صرفاً بر پایه‌ی نوآوری الگوریتمی یافته‌ای که پیامدهای مستقیمی برای طراحی سامانه‌های پیش‌بینی عملیاتی و حکمرانی آبی در مناطق آسیب‌پذیر دارد.

استناد: معروفی‌نیا، ا.، عصمت ساعتلو، س.م.، شکفتی، ص. و یزدانی‌مهر، ر. (۱۴۰۳). ارزیابی عملکرد روش‌های مبتنی بر محاسبات نرم برای پیش‌بینی روزانه جریان رودخانه با تفکیک زمانی روزانه. *اقلیم و بوم‌سازگان مناطق خشک و نیمه‌خشک*، ۲ (۲)، ۲۰۶-۲۲۶.

ناشر: دانشگاه سمنان

DOI: <https://doi.org/10.22075/ceasr.2025.40107.1063>

۱- مقدمه

پیش‌بینی دقیق رواناب روزانه رودخانه برای مدیریت مؤثر منابع آب ضروری است. این پیش‌بینی‌ها نقشی اساسی در فرایند تصمیم‌گیری مدیران منابع آب و سیاست‌گذاران ایفا می‌کنند و بر تخصیص آب، بهره‌برداری از مخازن، اقدامات کنترل سیلاب و راهبردهای کاهش خشکسالی تأثیر می‌گذارند (Kumar و همکاران، ۲۰۲۳). برای دستیابی به پیش‌بینی‌های قابل‌اطمینان، مدل‌های متعدد هیدرولوژیک توسعه داده شده و برای شبیه‌سازی رواناب رودخانه‌ای مورد استفاده قرار گرفته‌اند که می‌توان آن‌ها را به سه دسته تجربی، مفهومی و مبتنی بر فیزیک طبقه‌بندی کرد (Kim و Cho، ۲۰۲۲). مدل‌های تجربی ساده و داده‌محورند ولی تعمیم‌پذیری پایینی دارند؛ مدل‌های مفهومی تعادل بین پیچیدگی و دقت را فراهم می‌کنند؛ در حالی که مدل‌های مبتنی بر فیزیک با در نظر گرفتن فرایندهای واقعی و پارامترهای فضایی، دقت و قابلیت انتقال بالاتری برای پیش‌بینی جریان رودخانه ارائه می‌دهند (Adnan و همکاران، ۲۰۲۲). با وجود مزیت‌های مختلف مدل‌های مبتنی بر فیزیک علیرغم دقت بالا، به داده‌های ورودی فراوان، پارامترهای متعدد و زمان محاسباتی طولانی نیاز دارند و در شرایط کمبود داده یا عدم قطعیت پارامترها، عملکرد آن‌ها دچار ضعف می‌شود (Naabil و همکاران، ۲۰۱۷). در مقابل، مدل‌های یادگیری ماشین، با

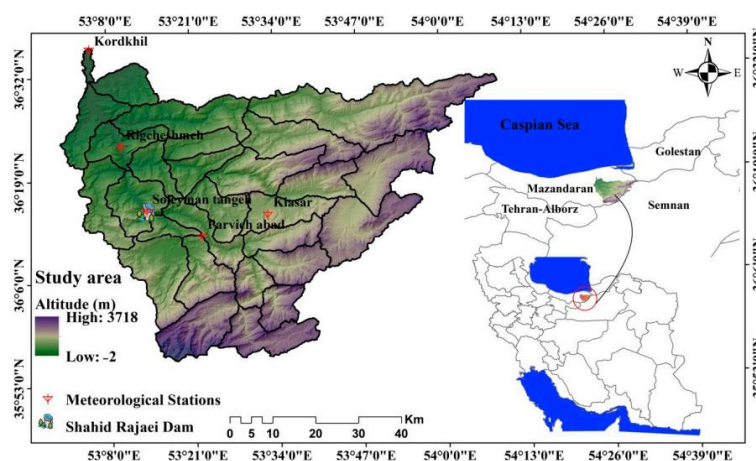
حداقل فرضیه فیزیکی، قادرند الگوهای پیچیده غیرخطی در داده‌های تاریخی را یاد بگیرند، نیاز به داده و محاسبات کمتری دارند و در بسیاری از موارد دقت پیش‌بینی بالاتری ارائه می‌دهند؛ هرچند تفسیرپذیری و تعمیم‌پذیری آن‌ها در شرایط تغییرات آب‌وهوایی یا حوضه‌های فاقد داده، چالش‌برانگیز است (Wang و همکاران، ۲۰۲۴). کاربرد روش‌های یادگیری ماشین و یادگیری عمیق از اواسط قرن بیستم در علوم و مهندسی محبوبیت گسترده‌ای یافته است (Meshram و همکاران، ۲۰۲۲). Kedam و همکاران (۲۰۲۴) به پیش‌بینی جریان رودخانه در حوضه نارمادا با استفاده از مدل‌های یادگیری ماشین پرداختند. در این پژوهش، مدل‌های CatBoost، LGBM، جنگل تصادفی و XGBoost روی داده‌های دوره ۱۹۷۸-۲۰۲۰ ارزیابی شدند. یافته‌ها نشان داد که جنگل تصادفی عملکرد برتری در پیش‌بینی جریان رودخانه دارد. Solanki و همکاران (۲۰۲۵) به بهبود پیش‌بینی جریان رودخانه با استفاده از مدل‌های هیدرولوژیک چندگانه و روش‌های یادگیری ماشین پرداختند. در پژوهش مذکور از مدل‌های هیدرولوژیک چندگانه و روش‌های یادگیری ماشین (رگرسیون خطی چندگانه، جنگل تصادفی، XGBoost و حافظه طولانی کوتاه‌مدت) برای پس‌پردازش خروجی مدل‌های هیدرولوژیک استفاده گردید. یافته‌های پژوهش مذکور نشان داد که ترکیب مدل‌های فیزیکی، داده‌های مشاهداتی هواشناسی و روش‌های یادگیری ماشین منجر به بهبود قابل توجهی در پیش‌بینی جریان رودخانه، به‌ویژه در تعیین شدت، زمان‌بندی و گستره سیل می‌شود (Solanki و همکاران، ۲۰۲۵). Zhou و همکاران (۲۰۲۵) به پیش‌بینی جریان رودخانه با ادغام محدودیت‌های هیدرولوژیک در مدل‌های یادگیری عمیق پرداختند. در پژوهش مذکور، یک مدل شبکه کانولوشنی گرافی تحت محدودیت هیدرولوژیک (HCGCN) توسعه داده شد که با در نظر گرفتن زمان تأخیر آب و جهت جریان، وابستگی‌های فضا-زمانی را به‌صورت تطبیقی یاد می‌گیرد. یافته‌های پژوهش مذکور نشان داد که HCGCN دقت پیش‌بینی جریان رودخانه را به‌طور چشمگیری بهبود می‌بخشد (Zhou و همکاران، ۲۰۲۵). Vinokić و همکاران (۲۰۲۵) به ارزیابی سه مدل یادگیری ماشین (TKAN، LSTM و TCN) برای پیش‌بینی جریان روزانه رودخانه و سنجش عدم قطعیت پرداختند. یافته‌های پژوهش نشان داد که TKAN عملکردی برتر از LSTM و نزدیک به TCN دارد، به‌طوری که در افق پیش‌بینی سه‌روزه، خطای مطلق میانگین و کارایی نش-ساتکلیف آن به‌ترتیب ۵/۷۹۹ متر مکعب بر ثانیه و ۰/۹۵۸ بود. همچنین، تحلیل عدم قطعیت و تفسیرپذیری مبتنی بر SHAP نشان داد که TKAN از پایداری قابل‌قبولی در پیش‌بینی‌های کوتاه‌مدت برخوردار است (Vinokić و همکاران، ۲۰۲۵). Grey و همکاران (۲۰۲۵) با تلفیق یادگیری ماشین و آمار مکانی، پیش‌بینی جریان روزانه در حوضه‌های فاقد ایستگاه را بهبود بخشیدند. در پژوهش مذکور، دو مدل LSTM و شبکه جریان مکانی (SSN) ارزیابی گردید. یافته‌های آن‌ها نشان داد که LSTM در اکثر موارد دارای عملکرد بهتری است. López-Chacón و همکاران (۲۰۲۵) با ترکیب یک مدل مبتنی بر فیزیک و یک الگوریتم یادگیری ماشین (جنگل تصادفی) مدلی ترکیبی برای بهبود پیش‌بینی دبی‌های بالا ارائه کردند. یافته‌ها نشان داد که مدل ترکیبی در مقایسه با مدل یادگیری ماشین منفرد، خطای جذر میانگین مربعات را ۲۳ تا ۳۳ درصد و بایاس را بیش از ۷۰ درصد برای دبی‌های بالاتر از دوره بازگشت ۱/۵ تا ۳ ساله کاهش می‌دهد (López-Chacón و همکاران، ۲۰۲۵). Nasir و همکاران (۲۰۲۵) به ارزیابی تطبیقی مدل‌های یادگیری ماشین برای پیش‌بینی جریان روزانه در حوضه بالادست رودخانه کلانگ در مالزی پرداختند. پیش‌پردازش داده‌ها با استفاده از هموارسازی کالمن و انتخاب ورودی‌های تأخیری بر اساس تحلیل خودهمبستگی انجام شد و مدل رگرسیون فرایند گوسی (GPR) با هسته نمایی به‌عنوان برترین مدل شناسایی گردید. یافته‌ها نشان داد که GPR با در نظر گرفتن وابستگی‌های زمانی، توانایی بالایی در بازتولید الگوهای پیچیده دارد (Nasir و همکاران، ۲۰۲۵). Danesh و همکاران (۲۰۲۵) به ارزیابی تطبیقی مدل‌های یادگیری ماشین و یادگیری عمیق برای پیش‌بینی جریان روزانه رودخانه در دو ایستگاه حوضه رودخانه مکنزی در ایالات متحده پرداختند. در پژوهش مذکور، سه مدل یادگیری ماشین (جنگل تصادفی، درخت تصمیم و K- نزدیک‌ترین همسایه) و سه مدل یادگیری عمیق (CNN، LSTM و مدل ترکیبی CNN-LSTM) مورد آزمون قرار گرفتند. یافته‌ها نشان داد که مدل‌های یادگیری عمیق، به‌ویژه

CNN-LSTM، عملکرد برتری دارند و در هر دو ایستگاه بهترین دقت را با مقادیر R^2 بیشتر از ۰/۸۸ و KGE بیشتر از ۰/۹۳ ارائه می‌دهند (Danesh و همکاران، ۲۰۲۵). اگرچه مدل‌های یادگیری ماشین و عمیق در پیش‌بینی جریان رودخانه پیشرفت چشمگیری داشته‌اند، شکاف‌هایی همچنان پابرجاست: بیشتر مدل‌ها فاقد ادغام منسجم دانش فیزیکی حوضه آبخیز، محدود به شرایط خاص جغرافیایی و آب‌وهوایی بوده و در پیش‌بینی دبی‌های اوج حیاتی برای مدیریت سیلاب دقت کافی ندارند. علاوه بر این، تعمیم‌پذیری در حوضه‌های فاقد داده و شفاف‌سازی عدم قطعیت و تفسیرپذیری مدل‌های پیچیده نیز کمتر مورد توجه قرار گرفته است. هدف این پژوهش، توسعه یک چارچوب هیبریدی نوآورانه است که با تلفیق محدودیت‌های فیزیکی، ویژگی‌های فضا-زمانی جریان و قابلیت‌های تطبیقی یادگیری عمیق، دقت پیش‌بینی دبی اوج را افزایش داده و خروجی‌هایی قابل اعتماد، قابل تفسیر و همراه با سنجش عدم قطعیت ارائه کند تا تصمیم‌گیری‌های مؤثر در مدیریت منابع آب و کاهش خطر سیلاب تسهیل گردد.

۲- مواد و روش‌ها

۲-۱- منطقه مورد مطالعه

تجن یکی از رودخانه‌های مستقل دریای خزر در زیرحوضه نکاء- تجن در شرق شهرستان ساری و جنوب شهرستان‌های نکاء و بهشهر جریان دارد و از رودخانه‌های مهم و دائمی استان مازندران است. حوضه آبخیز تجن با مساحت ۳۸۱۰ کیلومتر مربع در استان مازندران واقع شده و از جنوب توسط کوه‌های البرز و از شمال توسط دریای خزر محصور گردیده است. این حوضه بین طول‌های جغرافیایی $36^{\circ}53'N$ تا $36^{\circ}55'N$ شرقی و عرض‌های جغرافیایی $53^{\circ}36'N$ تا $53^{\circ}39'N$ شمالی قرار دارد. میانگین دمای سالانه در این منطقه حدود ۱۵ درجه سانتی‌گراد بوده و دارای اقلیم گرم و مرطوب با بارش سالانه‌ای در حدود ۷۶۰ میلی‌متر است (Saeidi و همکاران، ۲۰۲۲). کمترین و بیشترین ارتفاعات حوضه تجن به‌ترتیب ۲۶ و ۳۷۲۸ متر است. این رودخانه محل انجام فعالیت‌ها و عملیات گوناگون کشاورزی، آبی‌پروری و صنعتی از جمله سدسازی و استخراج شن و ماسه است، به‌طوری که میزان متوسط TDS اندازه‌گیری شده به‌طور مستقیم تحت تأثیر این فرایندها قرار دارد (Sun و همکاران، ۲۰۲۱). شکل ۱، موقعیت جغرافیایی حوضه آبخیز تجن را نشان می‌دهد.



شکل ۱. موقعیت جغرافیایی حوضه آبخیز تجن (Avand et al., 2024)

Fig. 1 Geographic location of the Tajan watershed (Avand et al., 2024)

۲-۲- داده‌های مطالعه

در پیش‌بینی جریان رودخانه، شناسایی و به‌کارگیری پارامترهای مؤثر نقش حیاتی در بهبود دقت مدل‌های هیدرولوژیک دارد. مهم‌ترین متغیر ورودی، بارش است که به‌عنوان محرک اصلی فرایند رواناب عمل می‌کند. دما و تبخیر-تعرق نیز بر میزان رطوبت خاک و در نتیجه بر تولید رواناب اثر می‌گذارند. رطوبت اولیه خاک و ویژگی‌های فیزیوگرافی حوضه (مانند شیب، ارتفاع، مساحت، کاربری اراضی و نوع خاک) رفتار هیدرولوژیک رودخانه را تعیین می‌کنند. دبی روز قبل یا مقادیر تأخیردار جریان نیز برای مدل‌های داده‌محور اهمیت دارند، زیرا رفتار جریان اغلب وابستگی زمانی دارد. در این پژوهش، به علت اثرات ناچیز پارامترهای هواشناسی (مانند رطوبت نسبی، سرعت باد و فشار هوا) از اثرات آنها صرف نظر گردید. لذا صرفاً از داده‌های بارش روزانه (بر حسب میلی‌متر)، تبخیر روزانه (بر حسب میلی‌متر) و دبی جریان رودخانه (متر مکعب بر ثانیه) با یک تا سه تأخیر به‌عنوان متغیرهای ورودی و دبی جریان رودخانه به‌عنوان متغیر خروجی استفاده گردید. سری زمانی داده‌ها به‌صورت روزانه بوده که اطلاعات ذکر شده در دوره زمانی ۱۳۴۸ لغایت پایان ۱۳۹۷ از شرکت آب منطقه‌ای مازندران اخذ گردید. تعداد کل نمونه‌های تحقیق شامل ۱۵۱۵۲۵ نمونه می‌باشد که از ۷۰ درصد داده‌ها (۱۰۶,۰۶۷ نمونه) در فرایند آموزش و ۳۰ درصد (۴۵,۴۵۸ نمونه) برای آزمون استفاده گردید.

۲-۳- مدل‌های پژوهش

مدل‌های یادگیری ماشین و عمیق، هر دو در پیش‌بینی هیدرولوژیک کاربرد دارند، اما عملکرد آنها به شدت به کیفیت داده و ساختار ورودی بستگی دارد. در شرایط داده‌های محدود و پرتنش (همانند حوضه‌های با نوسانات شدید جریان) مدل‌های کرنلی اغلب به دلیل تنظیم ذاتی پیچیدگی، پایدارتر و کم‌سوگیرتر عمل می‌کنند. در مقابل، مدل‌های عمیق با وجود ظرفیت بالا در یادگیری روابط غیرخطی، در صورت نبود داده‌های گسترده و مناسب، مستعد بیش‌برازش و عدم تعمیم‌پذیری هستند. در این پژوهش از یک مدل یادگیری ماشین (SVR) و دو مدل یادگیری عمیق (GAN و CNN) استفاده گردید.

۲-۳-۱- مدل شبکه مولد رقابتی (GAN)

شبکه‌های مولد رقابتی یا تخصصی (GANs) نشان‌دهنده پیشرفتی کلیدی در هوش مصنوعی هستند که چارچوبی قوی برای تولید داده‌های مصنوعی فراهم می‌کنند که به‌شدت به اطلاعات دنیای واقعی شبیه است (Goodfellow و همکاران، ۲۰۲۰). شبکه‌های GANs متشکل از دو شبکه عصبی مرتبط مولد و تشخیص‌دهنده (تمیزدهنده) هستند که در یک فرایند رقابتی پویا با یکدیگر تعامل دارند و چشم‌انداز مدل‌سازی تولیدی عمیق را بازتعریف می‌کنند (Krichen, ۲۰۲۳). با هدایت این تعامل، شبکه‌های GANs مرزهای تولید داده را در حوضه‌های گوناگون از ایجاد تصاویر تا تولید زبان فراتر می‌برند و تأثیر عمیقی در بازتعریف چگونگی درک و بازتولید توزیع‌های پیچیده داده‌ای توسط ماشین‌ها دارند. در این الگوریتم، مدل سازنده بر مبنای رقابت پیوسته و همزمان شبکه مولد (G) و تشخیص‌دهنده (D) تخمین زده می‌شود. برای یادگیری توزیع مولد P_g بر داده‌های x یک نوفه اولیه بر داده‌های ورودی $p_z(z)$ تعریف می‌شود. سپس یک نگاهت به فضای داده‌ها به صورت $G(z; \theta_g)$ که G یک تابع تمایز است که توسط یک پرسپترون چندلایه با مقدار θ_g تعریف می‌شود که خروجی آن به‌صورت عددی می‌باشد. این الگوریتم آموزش می‌بیند تا احتمال تخصیص برچسب صحیح توسط تمیزدهنده روی داده‌های ورودی واقعی و داده‌های ساخته شده با مولد افزایش یابد. در همین حال، مولد آموزش می‌بیند تا مقدار $\log(1 - D(G(z)))$ کمترین مقدار ممکن برسد. به بیان دیگر، مولد و تمیزدهنده بازی حداقل-حداکثر دو نفره زیر را با تابع مقدار $V(G, D)$ انجام می‌دهند که با استفاده از مقدار مورد انتظار (E) تعریف نمود. روابط ریاضی این مدل از (۱) تا (۴) به‌صورت زیر تعریف می‌گردد:

$$\min_G \max_D V(D, G) = E_{x \sim P_{data}(x)} [\log D(x)] + E_{z \sim P_z(z)} [\log(1 - D(G(z)))] \quad (1)$$

$$D_G^*(x) = \frac{P_{data}(x)}{P_{data}(x) + P_g(x)} \quad (2)$$

$$V(G, D) = \int P_{data}(x) \log(D(x)) dx + \int_z P_z(z) \log(1 - D(g(z))) \quad (3)$$

$$= \int_x P_{data}(x) \log(D(x)) + P_g(x) \log(1 - D(x)) dx$$

$$C(G) = \max_D V(G, D) = E_{x \sim P_{data}} [\log D_G^*(x)] + E_{z \sim P_z} [\log(1 - D_G^*(G(z)))] \quad (4)$$

$$C(G) = E_{x \sim P_{data}} [\log D_G^*(x)] + E_{x \sim P_g} [\log(1 - D_G^*(x))]$$

$$C(G) = E_{x \sim P_{data}} \left[\log \frac{P_{data}(x)}{P_{data}(x) + P_g(x)} \right] + E_{x \sim P_g} \left[\log \frac{P_g(x)}{P_{data}(x) + P_g(x)} \right]$$

که در روابط بالا به ترتیب z : یک توزیع تصادفی به عنوان داده‌های ورودی به مولد، $G(z)$: نمونه ساخته شده با استفاده از ورودی تصادفی، P_{data} : توزیع احتمالی داده‌های واقعی، P_z : توزیع احتمالی نویز ورودی (معمولاً نرمال استاندارد)، P_g : توزیع احتمالی داده‌های تولیدشده توسط مولد، x : نمونه داده (واقعی یا مصنوعی)، D_G^* : تشخیص‌دهنده بهینه (در حالت تعادل)، $E[\cdot]$: امید ریاضی (میانگین انتظاری)، $C(G)$: عملکرد (هزینه/بازده) مدل مولد G ، $V(G, D)$: تابع بازی (تعامل بین مولد و تشخیص‌دهنده)، $\max_D$: نشان‌دهنده این است که تشخیص‌دهنده D سعی می‌کند تا حداکثر دقت را در تشخیص داده‌های واقعی از مصنوعی داشته باشد.

۲-۳-۲- شبکه عصبی پیچشی (CNN)

شبکه‌های عصبی پیچشی یا همگشتی (CNNs) شبکه‌های عصبی هستند که در حداقل یکی از لایه‌های خود، از عملگر کانولوشن به جای ضرب ماتریسی معمول استفاده می‌کنند. با امکان پذیرش ورودی‌ها و تولید خروجی‌های چندمتغیره، CNNs مزیت شبکه‌ی پرسپترون چندلایه (MLP) را در پیش‌بینی سری‌های زمانی دارند (Vatanchi و همکاران، ۲۰۲۳) در حقیقت مدل مذکور بدون نیاز به اینکه مدل مستقیماً از مشاهدات تأخیری یاد بگیرد، می‌تواند ارتباطات تابعی دلخواه اما پیچیده‌ای را نیز فرا بگیرد. در نتیجه، مدل CNN می‌تواند بازتاب‌های بسیار مرتبط با مسئله پیش‌بینی را از میان مجموعه‌ای گسترده از ورودی‌ها استخراج کند (Ismail-Fawaz و همکاران، ۲۰۲۲). سه لایه اصلی یک مدل CNN شامل لایه کانولوشنی، لایه تجمیع و لایه کاملاً متصل می‌باشد. تعداد لایه‌ها یا انواع لایه‌ها در یک مدل CNN ساده می‌تواند به صورت پویا تغییر کند. لایه کانولوشنی با استفاده از چندین کرنل کانولوشنی، نمایش‌های ویژگی‌ای از ورودی‌ها را یاد می‌گیرد (Miau و Hung، ۲۰۲۰). یکی از دلایل کلیدی استفاده از لایه کانولوشنی به جای لایه کاملاً متصل این است که لایه کاملاً متصل منجر به تعداد زیادی پارامتر می‌شود که نیازمند منابع محاسباتی فراوانی است. خروجی‌های لایه کانولوشنی پیش از ارسال به لایه بعدی، از طریق یک تابع فعال‌سازی غیرخطی پردازش می‌شوند (Kumshe و همکاران، ۲۰۲۴). هنگامی که ورودی یک سیگنال یک‌بعدی (D-1) باشد، لایه کانولوشنی h با استفاده از مجموعه‌ای از فیلترهای کوچک $k = 1, \dots, N_k$ با اندازه $(L \times 1)$ طبق رابطه (۵) ساخته می‌شود:

$$h_i^k = f \left(\sum_{l=1}^L w_l^k X_{i+l} + b^k \right) \quad (5)$$

که در آن انتخاب رایج برای $f(\cdot)$ تابع واحد خطی اصلاح شده (ReLU) است. این لایه‌ها می‌توانند با سایر معماری‌ها برای انجام وظایف پیچیده‌تری مانند LSTM یا GRU ترکیب شوند (Ghimire و همکاران، ۲۰۲۱).

۲-۳-۳- مدل ماشین بردار پشتیبان رگرسیون (SVR)

مدل ماشین بردار پشتیبان (SVM) که توسط Cortes و Vapnik (۱۹۹۵) ارائه شده، یکی از روش‌های مؤثر برای حل مسائل دسته‌بندی و رگرسیون محسوب می‌شود (Bargam و همکاران، ۲۰۲۴). کاربرد گسترده ماشین بردار به دلیل عملکرد کارآمد آن است. در حقیقت SVM با استفاده از یک تابع خطی، ورودی‌ها را به فضایی با ابعاد بالا تصویر می‌کند و سپس با به‌کارگیری یک ابرصفحه بهینه، مجموعه داده‌ها را بر اساس جابجایی از هم تفکیک می‌کند (Yu و همکاران، ۲۰۲۰). هدف رگرسیون بردار پشتیبان یافتن تابع رگرسیونی $f(x)$ است که بتواند به‌طور مناسب رابطه بین مجموعه داده ورودی داده شده $x = \{x_1, x_2, x_3, \dots, x_n\}$ و مقادیر هدف $y = \{y_1, y_2, y_3, \dots, y_n\}$ را طبق رابطه (۶) توصیف کند:

$$f(x) = w\phi(x) + b \quad (6)$$

که در آن w بردار وزن تابع، b : عامل جابجایی و ϕ : نگاشت غیرخطی است که به‌منظور جلوگیری از برازش بیش از حد داده‌های واسنجی به کار می‌رود. در حقیقت، SVR از یک تابع هدف و یک تابع زیان برای به‌دست‌آوردن پارامترهای رگرسیونی در معادله فوق به‌صورت زیر مطابق روابط (۷) تا (۸) استفاده می‌کند:

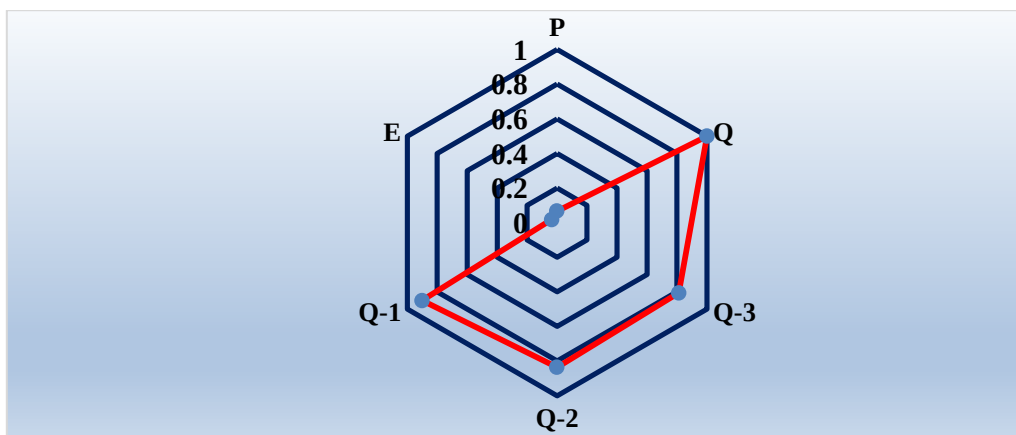
$$\min \frac{1}{2} \|w\|^2 + C \sum_1^n (\xi_i + \xi_i^*) \quad (7)$$

$$\text{subject to: } \begin{cases} y_i - (w\phi(x) + b) \leq \epsilon + \xi_i \\ (w\phi(x) + b) - y_i \leq \epsilon + \xi_i^* \\ \xi_i \geq 0, \xi_i^* \geq 0, i = 1, \dots, n \end{cases} \quad (8)$$

که در آن ϵ یک آستانه مثبت خطا است. همچنین C ضریب جریمه‌ای است که توسط کاربر تعیین می‌شود. همچنین ξ_i و ξ_i^* متغیرهای کمبود مثبتی هستند که برای تعیین انحراف داده‌های واسنجی از ϵ به‌کار می‌روند (Merufinia و همکاران، ۲۰۲۳).

۲-۴- سناریو و ترکیب مدل

در این پژوهش، از متغیرهای بارش (P_t)، تبخیر (E_t)، دبی با یک روز تأخیر (Q_{t-1})، دبی با دو روز تأخیر (Q_{t-2})، دبی با سه روز تأخیر (Q_{t-3}) و خروجی جریان نیز Q_t می‌باشد. جهت ارتباط همبستگی بین متغیرهای ورودی و خروجی از ضریب پیرسون استفاده شد و نتایج در نمودار راداری (عنکبوتی) در شکل ۲ نشان داده شده است. همچنین، ترکیب‌های مختلف و سناریوهای پیشنهادی در جدول ۱ نشان داده شده است.



شکل ۲. نمودار راداری (عنکبوتی) بین متغیرهای ورودی و خروجی

Fig. 2 Radar (spider) chart of relationships between input and output variables

جدول ۱. ترکیب و سناریوهای تحقیق

Table 1. Model configurations and research scenarios

سناریو Scenario	متغیرهای ورودی Input Variables	متغیر خروجی Output Variable
SN ₁	P(t)	Q(t)
SN ₂	E(t), P(t)	Q(t)
SN ₃	E(t), P(t), Q(t-3)	Q(t)
SN ₄	Q(t-1), Q(t-2), Q(t-3)	Q(t)
SN ₅	E(t), P(t), Q(t-1), Q(t-2), Q(t-3)	Q(t)

۵-۲- خلاصه پارامترهای آماری محدوده مطالعاتی

پارامترهای آماری (از جمله میانگین، انحراف معیار، ضریب تغییرات، چولگی و کشیدگی) اطلاعاتی کلیدی در باره رفتار هیدرولوژیک حوضه آبخیز فراهم می‌کنند. این شاخص‌ها ساختار آماری داده‌های جریان رودخانه را تبیین کرده و به شناسایی الگوهای فصلی، نوسانات شدید و روندهای بلندمدت کمک می‌کنند. وجود تغییرپذیری بالا یا چولگی قابل توجه در داده‌ها، نشان‌دهنده ناهمگونی رژیم جریان و چالش‌های پیش روی مدل‌های پیش‌بینی است. این پارامترها در طراحی و ارزیابی مدل‌های پیش‌بینی جریان رودخانه اهمیت زیادی دارند، زیرا امکان درک بهتر از پویایی سیستم هیدرولوژیک، تشخیص داده‌های پرت، و انتخاب ورودی‌های مناسب را فراهم می‌کنند. جدول ۲، خلاصه پارامترهای آماری محدوده مطالعاتی برای کل داده‌ها را نشان می‌دهد.

جدول ۲. خلاصه پارامترهای آماری محدوده مطالعاتی برای کل داده‌ها

Table 2. Summary of statistical parameters for the study area based on the entire dataset

پارامترهای آماری محدوده مطالعاتی برای کل داده‌ها							
Statistical parameters of the study area for the entire dataset							
ردیف Row	پارامتر Parameter	بیشینه Maximum	کمینه Minimum	میانگین Mean	انحراف استاندارد Standard Deviation	چولگی Skewness	کشیدگی Kurtosis
1	بارش (mm) Rainfall	240	0	1.887	6.733	9.387	151.903
2	تبخیر (mm/day) Evaporation	3.203	1/0	2.576	2.677	5.325	244.289
3	دبی (m ³ /s) Discharge	460	0	8.027	14.767	5.421	60.485

۲-۶- معیارهای ارزیابی مدل

در مدل‌سازی و پیش‌بینی جریان رودخانه، از معیارهای آماری متعددی برای سنجش دقت و کارایی مدل‌ها استفاده می‌شود. این معیارها نقشی اساسی در ارزیابی قابلیت پیش‌بینی مدل‌ها ایفا می‌کنند؛ زیرا نه تنها دقت مدل را کمی سنجیده، بلکه به محققان کمک می‌کنند تا نقاط قوت و ضعف روش‌های مختلف را شناسایی کرده، مدل مناسب‌تری را برای کاربردهای عملیاتی مانند مدیریت منابع آب، هشدار سیل و برنامه‌ریزی‌های بلندمدت انتخاب نمایند. جدول ۳ معیارهای مختلف ارزیابی، محدوده قابل تعریف و مقدار بهینه آن را نشان می‌دهد.

۳- نتایج و بحث

در این بخش، نتایج حاصل از پیاده‌سازی مدل‌های پیش‌بینی دبی رودخانه بر اساس متغیرهای هیدرومتئورولوژیک کلیدی (شامل بارش، تبخیر و مقادیر دبی با تأخیرهای مختلف) ارائه گردید. پیش از مدل‌سازی، ساختار آماری داده‌های بلندمدت (۱۳۴۸-۱۳۹۷) با استفاده از شاخص‌های توصیفی و همبستگی‌های خطی مورد ارزیابی قرار گرفت تا اطلاعات لازم برای انتخاب ورودی‌های بهینه فراهم شود. داده‌های جمع‌آوری شده پس از تقسیم‌بندی (به نسبت ۷۰:۳۰)، به ترتیب در فرایندهای آموزش و آزمون مورد استفاده قرار گرفتند.

جدول ۳. معیارهای ارزیابی مدل
Table 3. Model evaluation metrics

ردیف Row	نام رابطه Name of Metric	رابطه ریاضی Mathematical Formula	محدوده تعریف Defined Range	مقدار بهینه Optimal Value
1	ضریب تبیین (R^2) Coefficient of Determination	$R^2 = \frac{[\sum_{i=1}^n (Q_{obs,i} - \overline{Q_{obs}})(Q_{sim,i} - \overline{Q_{sim}})]^2}{\sum_{i=1}^n (Q_{obs,i} - \overline{Q_{obs}})^2 \sum_{i=1}^n (Q_{sim,i} - \overline{Q_{sim}})^2}$	$(0, 1]$	1
2	میانگین قدر مطلق خطا (MAE) Mean Absolute Error	$MAE = \frac{1}{n} \sum_{i=1}^n Q_{sim,i} - Q_{obs,i} $	$[0, +\infty)$	0
3	ریشه میانگین مربعات خطا (RMSE) Root Mean Square Error	$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (Q_{sim,i} - Q_{obs,i})^2}$	$[0, +\infty)$	0
4	درصد سوگیری یا ناریبی (PBIAS) Percentage Bias	$PBIAS = 100 \times \frac{\sum_{i=1}^n (Q_{obs,i} - Q_{sim,i})}{\sum_{i=1}^n Q_{obs,i}}$	$(-\infty, +\infty)$	0
5	میانگین درصد خطای مطلق (MAPE) Mean Absolute Percentage Error	$MAPE = \frac{100}{n} \sum_{i=1}^n \left \frac{Q_{obs,i} - Q_{sim,i}}{Q_{obs,i}} \right $	$[0, +\infty)$	0
6	ضریب کلینگ-گوپتا (KGE) Kling-Gupta Efficiency	$KGE = 1 - \sqrt{(r-1)^2 + (\alpha-1)^2 + (\beta-1)^2}$	$(-\infty, 1]$	1

سپس، عملیات پیش‌پردازش و مدیریت داده‌ها (از قبیل آزمون‌های همگنی، شناسایی و حذف داده‌های پرت، بازسازی داده‌های مفقوده و نرمال‌سازی داده) انجام گرفت. نتایج جدول (۲) بر اساس شاخص‌های آماری نشان داد که در پارامتر بارش، میانگین اندک ۱/۸۸۷ میلی‌متر در مقابل بیشینه ۲۴۰ میلی‌متری و انحراف استاندارد ۶/۷۳۳، بیانگر ضریب تغییرات بسیار بالا و وقوع بارش‌های سیل‌آسا در بازه‌های زمانی کوتاه است که این امر با مقدار چولگی مثبت ۹/۳۸۷ و کشیدگی قابل توجه ۱۵۱/۹۰۳ تأیید می‌شود. وضعیت مشابهی در داده‌های دبی مشاهده می‌گردد، به‌طوری که میانگین ۸/۰۲۷ مترمکعب بر ثانیه با انحراف استاندارد ۱۴/۷۶۷ همراه شده است که نشان می‌دهد تغییرات جریان بیش از ۱/۸ برابر میانگین بوده و سیستم هیدرولوژیکی منطقه به‌شدت تحت تأثیر رویدادهای حدی قرار دارد. همچنین پارامتر تبخیر با وجود میانگین ۲/۵۷۶ میلی‌متر در روز، بیشترین میزان کشیدگی را با مقدار ۲۴۴/۲۸۹ به خود اختصاص داده است که حاکی از وجود روزهایی با تبخیر بسیار شدید و خارج از الگوی معمول منطقه می‌باشد. در مجموع، بالا بودن شاخص‌های گشتاور سوم و چهارم در هر سه متغیر نشان می‌دهد که رفتار هیدرواقليمی محدوده مطالعاتی از الگوهای خطی و نرمال پیروی نکرده و تحلیل‌های آتی باید بر مبنای مدل‌های سازگار با داده‌های چوله و با دنباله بلند انجام گیرد. پنج ترکیب مدل و سناریو در این پژوهش مورد بررسی قرار گرفت تا اثرات ترکیب‌های مدل در میزان دقت مورد ارزیابی قرار گیرد. در سناریوهای اول (SN_1) و دوم (SN_2)، دبی پایه رودخانه (یعنی مقادیر دبی جریان رودخانه با تأخیرهای یک تا سه روز) به‌عنوان ورودی مدل‌ها در نظر گرفته نشده است و تنها از متغیرهای هواشناسی مانند بارش و دما استفاده شده است تا ارزیابی گردد که در غیاب دبی پایه، مدل تا چه حد قادر

به پیش‌بینی جریان خواهد شد. نتایج جدول ۴ به‌وضوح نشان می‌دهد که هر سه مدل CNN، GAN و SVR در این دو سناریو عملکرد بسیار ضعیفی داشته به طوری که ضریب R^2 در محدوده ۰/۰۰۵ تا ۰/۰۸۴ و خطای RMSE در محدوده ۱۴-۱۵ متر مکعب بر ثانیه است. در حالی که در سناریوهای بعدی (به‌ویژه SN_4 و SN_5) این خطا به‌طور چشمگیری به کمتر از ۷ متر مکعب بر ثانیه کاهش یافته است. از دیدگاه کمی، ضعف عملکرد در SN_1 و SN_2 به دلیل عدم در نظر گرفتن پایداری ذاتی سری زمانی جریان رودخانه است. دبی رودخانه دارای همبستگی زمانی قوی است؛ یعنی مقادیر جریان در یک زمان به‌طور مستقیم تحت تأثیر مقادیر قبلی همان سری قرار دارد. حذف این اطلاعات اساسی، مدل‌ها را مجبور می‌کند تا تنها از سیگنال‌های ضعیف و غیرمستقیم (مانند بارش روزانه) برای پیش‌بینی جریان استفاده کنند؛ در حالی که پاسخ هیدرولوژیک حوضه به بارش به دلیل تأخیرهای هیدرولوژیک پیچیده و غیرخطی است. لذا نتایج دو سناریوی اول نشان داد که ضعف عملکرد پیش‌بینی، ریشه در ساختار ورودی مدل، نه در خود الگوریتم‌ها، دارد. از دیدگاه کمی، در مدل CNN، سناریوی SN_5 به‌عنوان بهترین حالت شناسایی می‌شود که در فاز آزمون دارای R^2 برابر با ۰/۸۱۸، RMSE معادل ۶/۵۸۷ متر مکعب بر ثانیه، KGE برابر با ۰/۸۶۲ و PBIAS معادل ۹/۸۶۷ درصد است. در مقابل، سناریوهای SN_1 و SN_2 عملکرد بسیار ضعیفی داشته‌اند که در آن‌ها R^2 به تقریب صفر و RMSE به‌طور میانگین بیش از ۱۴/۸ متر مکعب بر ثانیه رسیده است. در مدل GAN نیز بهترین عملکرد مربوط به سناریوی SN_5 است که در فاز آزمون R^2 معادل ۰/۸۱۳، RMSE برابر با ۶/۵۸۶ متر مکعب بر ثانیه، KGE معادل ۰/۷۸۴ و PBIAS مثبت و بسیار نزدیک به صفر دارد؛ در حالی که بدترین عملکرد در همین مدل متعلق به سناریوهای SN_1 و SN_2 است که R^2 آن‌ها به ترتیب ۰/۰۱۸ و ۰/۰۳۵ و RMSE آن‌ها در محدوده ۱۴/۷ تا ۴/۸ متر مکعب بر ثانیه قرار دارد. در مدل SVR، سناریوی SN_5 به‌وضوح بهترین نتایج را ارائه می‌دهد. در فاز آزمون این سناریو دارای R^2 برابر با ۰/۸۵، RMSE معادل ۵/۶۷۵ متر مکعب بر ثانیه، KGE برابر با ۰/۸۷۷ و PBIAS بسیار نزدیک به صفر (۰/۴۷۵ درصد) است. در حالی که ضعیف‌ترین عملکرد در این مدل نیز مربوط به SN_1 است که R^2 آن ۰/۰۱۲ و RMSE آن ۱۴/۵۹۶ متر مکعب بر ثانیه است. در هر سه مدل، SN_1 و SN_2 عملکرد نامناسبی دارند، در حالی که SN_4 و SN_5 بالاترین سطح همبستگی و کمترین خطا را نشان می‌دهند، که گویای یادگیری پایدار و تعمیم‌پذیری مناسب ساختار ورودی این سناریوهاست. با مقایسه عملکرد سه مدل در سناریوهای بهینه‌شان در فاز آزمون، مدل SVR به‌عنوان برترین مدل شناسایی می‌شود، زیرا نه تنها کمترین مقدار RMSE (۵/۶۷۵ متر مکعب بر ثانیه) و بیشترین KGE (۰/۸۷۷) را در میان تمام ترکیب‌های مدل-سناریو ارائه می‌دهد، بلکه PBIAS آن نیز بسیار نزدیک به صفر است که نشان‌دهنده عدم سوگرایی سیستماتیک است. در مقابل، مدل CNN در بهترین حالت خود (SN_5) دارای RMSE بالاتری (۶/۵۸۷ متر مکعب بر ثانیه) و PBIAS بزرگ‌تر (۹/۸۶۷ درصد) است، که نشان‌دهنده کم‌برآوردی یا بیش‌برآوردی سیستماتیک نسبتاً قابل توجه است. از این رو، SVR به‌عنوان مدل برتر و CNN به‌عنوان ضعیف‌ترین مدل در این مطالعه در نظر گرفته می‌شود، هرچند این تفاوت به‌مراتب کمتر از شکاف عملکردی بین سناریوهای ضعیف (SN_1 و SN_2) و سناریوهای برتر (SN_4 و SN_5) است. بنابراین، مدل SVR در سناریوی SN_5 (بهترین مدل) در فاز آموزش ۷۰/۶ درصد و در فاز آزمون ۶۱/۷ درصد عملکرد بهتری نسبت به بدترین نتایج که مربوط به مدل CNN در سناریوی SN_1 است دارد. این بهبود چشمگیر نشان‌دهنده تأثیر همزمان استفاده از یک مدل کارآمد (SVR) و یک ساختار ورودی غنی (شامل دبی‌های تأخیری و متغیرهای هواشناسی مرتبط) در مقابل یک مدل کم‌اثر با ورودی ناقص است. همچنین، کاهش درصد بهبود از فاز آموزش به فاز آزمون (از ۷۰/۶ به ۶۱/۷ درصد) گویای تعمیم‌پذیری خوب مدل برتر است، چرا که اختلاف عملکرد آن در دو فاز محدود بوده و دچار بیش‌برازش نشده است. بررسی مقایسه‌ای بین مدل یادگیری ماشین SVR و دو مدل یادگیری عمیق CNN و GAN در این مطالعه نشان می‌دهد که SVR عملکرد کلی برتری به‌ویژه در فاز آزمون از خود ارائه داده است. این برتری از دو دیدگاه الگوریتمی و داده‌محور قابل تبیین و تفسیر است. SVR به‌عنوان یک مدل کرنلی، با حداکثرسازی حاشیه بین نقاط پشتیبان، تمایل کمی به بیش‌برازش دارد

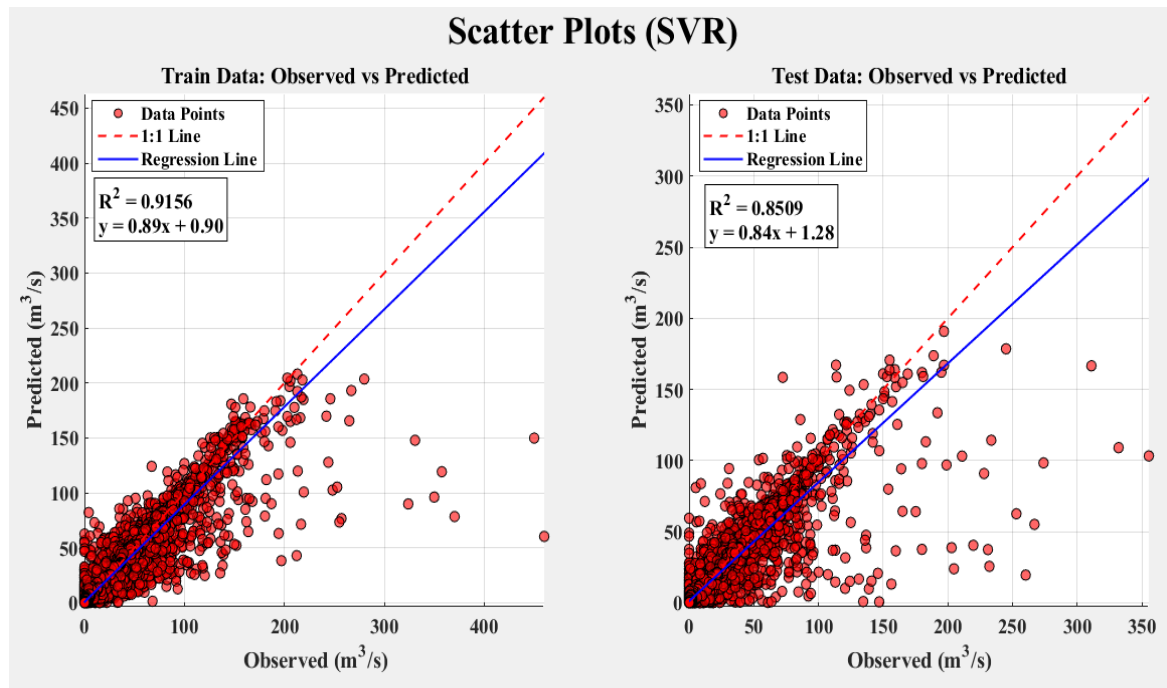
و در شرایط داده‌های محدود مانند سری‌های زمانی هیدرولوژیک کوتاه‌مدت تعادل مناسبی بین پیچیدگی و خطای آموزشی برقرار می‌کند. در مقابل، مدل‌های یادگیری عمیق مانند CNN و GAN به دلیل تعداد بالای پارامترها و نیاز به داده‌های گسترده، در این شرایط تمایل به یادگیری نویز یا ساختارهای بی‌ربط دارند که تعمیم‌پذیری آن‌ها را کاهش می‌دهد و در عملکرد ضعیف‌تر (RMSE بالاتر و PBIAS ناپایدار) منعکس شده است. علاوه بر این، ساختار زمانی پیوسته و فصلی داده‌های رودخانه‌ای با ماهیت SVR که با کرنل‌های غیرخطی می‌تواند روابط ورودی-خروجی را به‌صورت فشرده مدل کند سازگارتر است، در حالی که CNN برای سری‌های کوتاه کارایی محدودی دارد و GAN به‌عنوان مدلی برای تولید داده در مسائل رگرسیونی تک‌خروجی پایداری آموزشی کمتری از خود نشان می‌دهد. در مجموع، برتری SVR در این مطالعه نه نشان‌دهنده برتری ذاتی آن بر یادگیری عمیق، بلکه تناسب بالاتر آن با ماهیت داده‌ها و محدودیت‌های حجمی و ساختاری مسئله است. در شرایطی که داده‌های آموزشی محدود هستند و روابط ورودی-خروجی اگرچه غیرخطی اما نسبتاً پیوسته و پایدار باشند، مدل‌های مبتنی بر کرنل مانند SVR اغلب عملکردی قوی‌تر، پایدارتر و تفسیرپذیرتر از مدل‌های شبکه‌های عمیق از خود نشان می‌دهند. ارزیابی تطبیقی مدل‌های داده‌محور برای پیش‌بینی روزانه‌ی دبی رودخانه در حوضه‌ی آبخیز تجن نشان می‌دهد که مدل SVR هنگامی که با ساختار بهینه‌ی ورودی (شامل مقادیر دبی با تأخیر و متغیرهای هواشناسی کلیدی) ترکیب شود، از نظر دقت پیش‌بینی، پایداری آماری، کنترل سوگیری و تعمیم‌پذیری، عملکردی برتر نسبت به مدل‌های CNN و GAN دارد. این مدل در فاز آزمون، ضریب R^2 برابر با ۰.۸۵۰، RMSE معادل ۵.۶۷۵ متر مکعب بر ثانیه، PBIAS نزدیک به صفر (۰.۴۷۵٪) و KGE برابر با ۰.۸۷۷ را به‌دست آورد. این یافته‌ها با نتایج مطالعات اخیر در حوضه‌ی مدل‌سازی هیدرولوژیک هم‌راستا است و آن‌ها را گسترش می‌دهد. به‌عنوان مثال، Wang و همکاران (۲۰۲۲) و Ikram و همکاران (۲۰۲۳) نشان دادند که مدل SVR به‌ویژه هنگامی که با تکنیک‌های تجزیه (مانند EEMD یا EWT) یا الگوریتم‌های بهینه‌سازی ترکیب شوند در پیش‌بینی دبی رودخانه در اقلیم‌های متنوع، عملکردی ثابت و دقیق ارائه می‌دهند (Wang و همکاران، ۲۰۲۲؛ Ikram و همکاران، ۲۰۲۳). به‌طور مشابه، Khosravi و همکاران (۲۰۲۲) در مطالعه‌ای در شمال ایران، عملکرد بهبودیافته‌ی SVR را هنگام بهینه‌سازی با الگوریتم خفاش (BAT) گزارش کردند، هرچند بهترین مقدار KGE آن‌ها (حدود ۰/۸۲-۰/۸۵) همچنان کمتر از مقدار ۰/۸۷۷ به‌دست آمده در این پژوهش بود؛ علیرغم مشابهت در شرایط منطقه‌ای. این امر نشان می‌دهد که طراحی ورودی، به‌ویژه گنجاندن صریح دبی‌های قبلی (Q_{t-1} , Q_{t-2} , Q_{t-3}) نقشی تعیین‌کننده‌تر از پیچیدگی الگوریتمی دارد (Khosravi و همکاران، ۲۰۲۲). در حقیقت، نتایج این پژوهش تأیید می‌کند که مدل‌هایی که از دبی‌های تأخیری محروم شده‌اند (سناریوهای ۱ و ۲)، در هر سه معماری، ضریب R^2 کمتر از ۰.۰۴ داشتند که بیانگر این واقعیت هیدرولوژیک است. پایداری ذاتی سری‌زمانی جریان، یک پیش‌بینی‌کننده‌ی غیرقابل چشم‌پوشی در مدل‌سازی داده‌محور است، که یافته‌ای که Meshram و همکاران (۲۰۲۲) و Ayana و همکاران (۲۰۲۳) نیز با تأکید بر ضرورت ادغام مؤلفه‌های تناوبی و خودهمبستگی زمانی، به آن اشاره کرده‌اند. قابل توجه است که عملکرد ضعیف‌تر مدل‌های یادگیری عمیق این پژوهش (CNN و GAN)، در ظاهر با ادعاهای برخی پژوهش‌های اخیر (مانند Xu و همکاران، ۲۰۲۲؛ Latif و Ahmed، ۲۰۲۳) که بر برتری ذاتی این معماری‌ها در مسائل هیدرولوژیک تأکید دارند، در تضاد است. با این حال، این تناقض ظاهری با در نظر گرفتن دو عامل کلیدی حجم داده و ساختار مسئله قابل تبیین است. در حقیقت آن‌ها از داده‌های گسترده‌ی ماهانه (چند دهه از ایستگاه‌های متعدد در سیستم‌های بزرگ رودخانه‌ای) و اغلب از شاخص‌های اقلیمی کمکی استفاده کرده‌اند، شرایطی که از بیش‌برازش در مدل‌های با پارامتر زیاد جلوگیری می‌کند. در مقابل، این پژوهش بر اساس داده‌های روزانه‌ی یک ایستگاه در حوضه‌ای نیمه‌خشک با نوسانات شدید انجام شده است. حجم نسبتاً محدود داده‌ها (حدود ۱۸,۰۰۰ رکورد روزانه در ۵۰ سال، که پس از پیش‌پردازش کاهش یافته است) ظرفیت مؤثر مدل‌های عمیق را محدود می‌کند همچنین، عدنان و همکاران (۲۰۲۲) و کومار و همکاران (۲۰۲۳) اشاره کرده‌اند حتی هنگام استفاده از مدل‌های

پیچیده‌ی ترکیبی مانند CatBoost یا ANFIS مزیت‌های حاشیه‌ای پیچیدگی الگوریتمی در شرایط کمبود داده، به سرعت اشباع می‌شود. در چنین شرایطی، مدل‌های مبتنی بر کرنل مانند SVR که ذاتاً از طریق بیشینه‌سازی حاشیه، پیچیدگی را تنظیم می‌کنند، در برابر نویز و شکاف داده‌ای مقاوم‌تر بوده و از تمایل به بیش‌برازش که در مدل‌های CNN و GAN این پژوهش مشهود بود در امان می‌مانند (Adnan et al., 2022; Kumar et al., 2023). علاوه بر این، PBIAS بسیار نزدیک به صفر مدل SVR (0.475٪) بر برجسته‌ترین مزیت عملیاتی آن یعنی تعادل در پیش‌بینی جریان‌های کم و زیاد تأکید دارد؛ ویژگی‌ای حیاتی در حوضه‌هایی با چولگی بالا (۸۰۲۷) و کشیدگی شدید (۶۰.۴۸۵)، که وجود سیلاب‌های نادر ولی شدید را نشان می‌دهند. این در حالی است که مدل CNN در بهترین حالت خود PBIAS (SN5) معادل ۹.۸۶۷٪ داشت که نشان‌دهنده‌ی کم‌برآوردی یا بیش‌برآوردی سیستماتیک در رویدادهای اوج است. مطالعاتی مانند Ayana و همکاران (۲۰۲۳) مقادیر بسیار بالای NSE (بزرگ‌تر از ۰/۹۷) را با استفاده از مدل‌های پیشرفته‌ی یادگیری عمیق (مانند BiGRU) گزارش کرده‌اند، چنین عملکردی اغلب ناشی از استفاده از تکنیک‌های هموارسازی (میانگین متحرک) یا داده‌های ماهانه که نوسانات تصادفی را کاهش می‌دهند است. در مقابل، این پژوهش با حفظ ساختار زمانی خام داده‌ها بدون هموارسازی مصنوعی واقع‌گرایی هیدرولوژیک مسئله‌ی پیش‌بینی را حفظ کرده است و نتایج آن برای سناریوهای عملیاتی جایی که تغییرات ناگهانی جریان باید در زمان واقعی پیش‌بینی شوند قابل تعمیم‌تر است (Ayana et al., 2023). در مجموع، برتری SVR در این مطالعه نه نشان‌دهنده‌ی محدودیت ذاتی یادگیری عمیق، بلکه گواهی بر اصلی کلیدی در هیدروانفورماتیک است: انتخاب مدل باید بر اساس ویژگی‌های داده و محدودیت‌های مسئله باشد، نه صرفاً بر پایه‌ی نوآوری الگوریتمی. در حوضه‌هایی با داده‌های روزانه‌ی محدود و ناپایدار و با وابستگی زمانی قوی، یک مدل SVR خوب ساخته‌شده با ورودی‌های فیزیکی محور، راه‌حلی دقیق‌تر، پایدارتر و تفسیرپذیرتر نسبت به جایگزین‌های پیچیده‌تر ارائه می‌دهد یافته‌ای که پیامدهای مستقیمی برای حکمرانی پایدار آب در مناطق کم‌داده و تحت تنش اقلیمی مانند شمال غرب ایران دارد.

جدول ۴. نتایج عملکرد مدل‌های پژوهش در دو فاز آموزش و آزمون

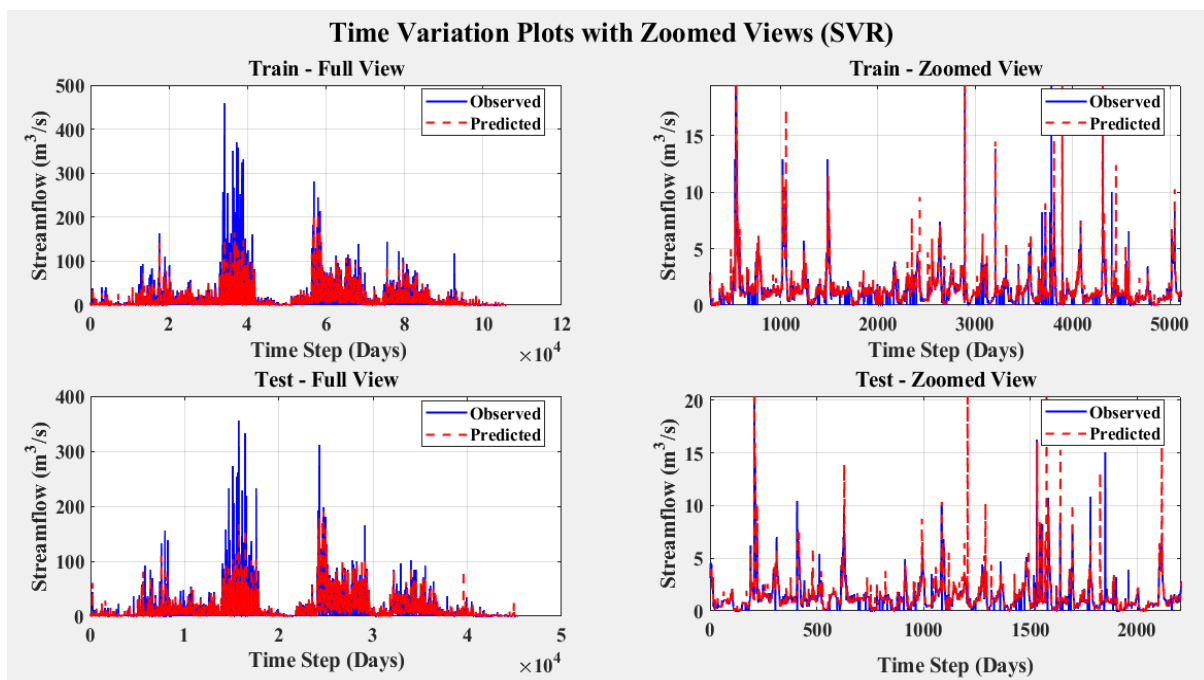
Table 4. Performance results of research models in two phases of training and testing

سناریوها Scenarios	فرایند آموزش (70%) Training Process				فرایند آزمون (30%) Testing Process			
	R ²	RMSE (m ³ /s)	PBIAS (%)	KGE (%)	R ²	RMSE (m ³ /s)	PBIAS (%)	KGE (%)
CNN model								
SN ₁	0.005	14.7	8.150	0.310	0.005	14.834	7.169	-0.310
SN ₂	0.011	14.603	6.501	0.242	0.0109	14.913	5.939	-0.247
SN ₃	0.804	7.288	17.255	0.773	0.774	7.746	16.91	0.784
SN ₄	0.848	5.983	-2.794	0.787	0.799	6.847	-0.3011	0.764
SN ₅	0.821	6.486	10.903	0.854	0.818	6.587	9.867	0.862
GAN model								
SN ₁	0.021	14.597	0.300	0.187	0.018	14.674	1.267	-0.200
SN ₂	0.037	14.389	4.189	0.146	0.035	14.748	4.701	-0.155
SN ₃	0.822	6.433	-5.752	0.796	0.847	5.717	-5.156	0.816
SN ₄	0.829	6.287	-4.349	0.795	0.823	6.270	-3.330	0.792
SN ₅	0.830	6.138	8.763	0.807	0.813	6.586	0.784	8.882
SVR model								
SN ₁	0.021	14.641	-0.190	0.207	0.012	14.596	0.129	0.231
SN ₂	0.0841	14.221	0.035	0.065	0.043	14.359	0.131	0.121
SN ₃	0.845	6.061	0.053	0.788	0.807	6.633	0.253	0.766
SN ₄	0.898	4.738	0.075	0.903	0.828	6.085	0.562	0.863
SN ₅	0.915	4.323	0.032	0.915	0.850	5.675	0.475	0.877



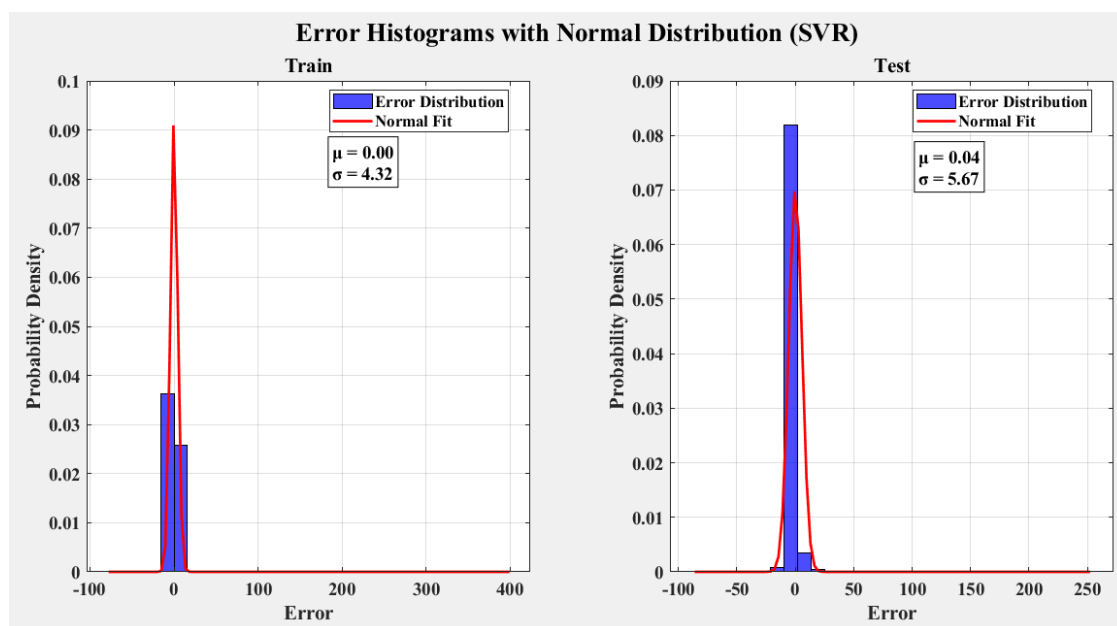
شکل ۳. نمودار پراکندگی بین داده‌های پیش‌بینی و مشاهداتی در مدل SVR

Fig. 3. Scatter plot between predicted and observed data of the SVR model



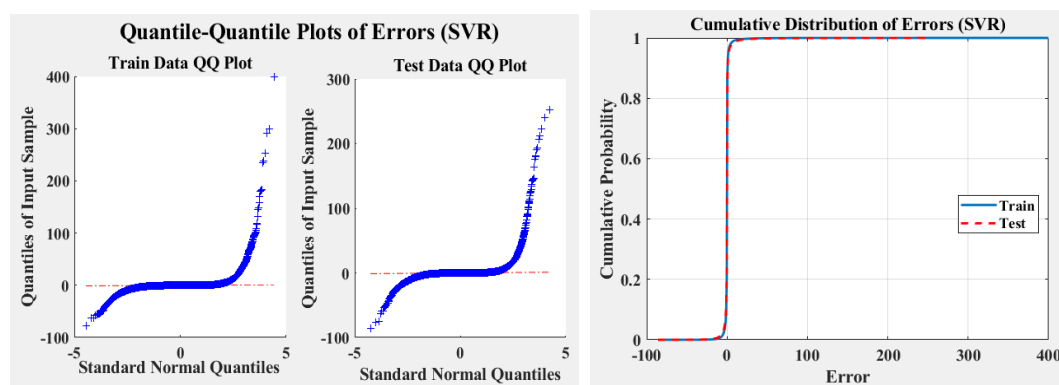
شکل ۴. نمودار سری زمانی بین داده‌های مشاهداتی و پیش‌بینی (SVR)

Fig. 4 Time series plot of observed and predicted data (SVR)



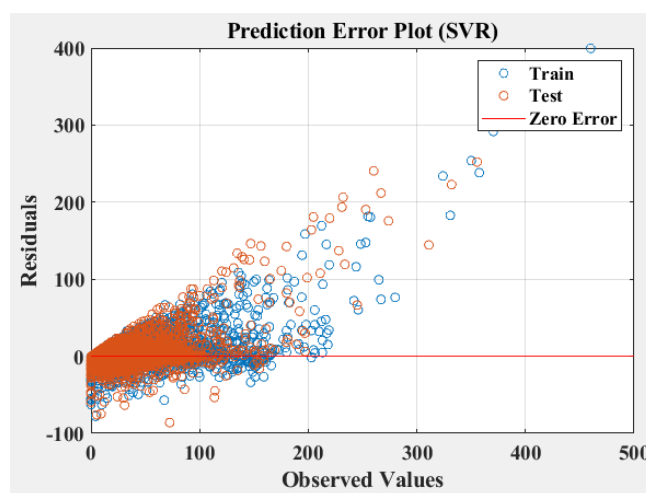
شکل ۵. نمودار هیستوگرام خطا و توزیع فراوانی مدل SVR

Fig. 5 Histogram of errors and frequency distribution of the SVR model



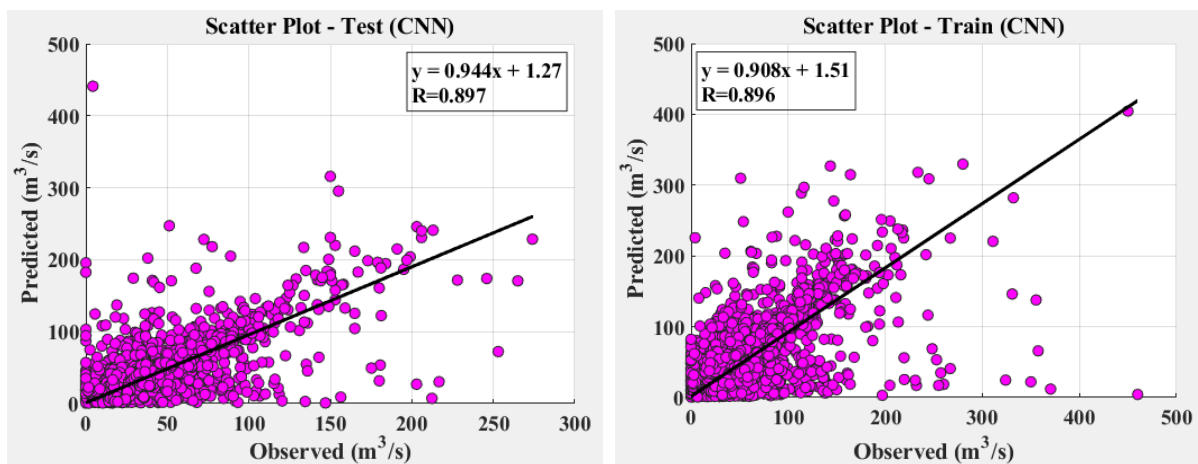
شکل ۶. نمودار توزیع تجمعی خطا و نمودار چندک-چندک خطاها در مدل SVR

Fig. 6 Cumulative error distribution plot and quantile-quantile (Q-Q) plot of errors for SVR model



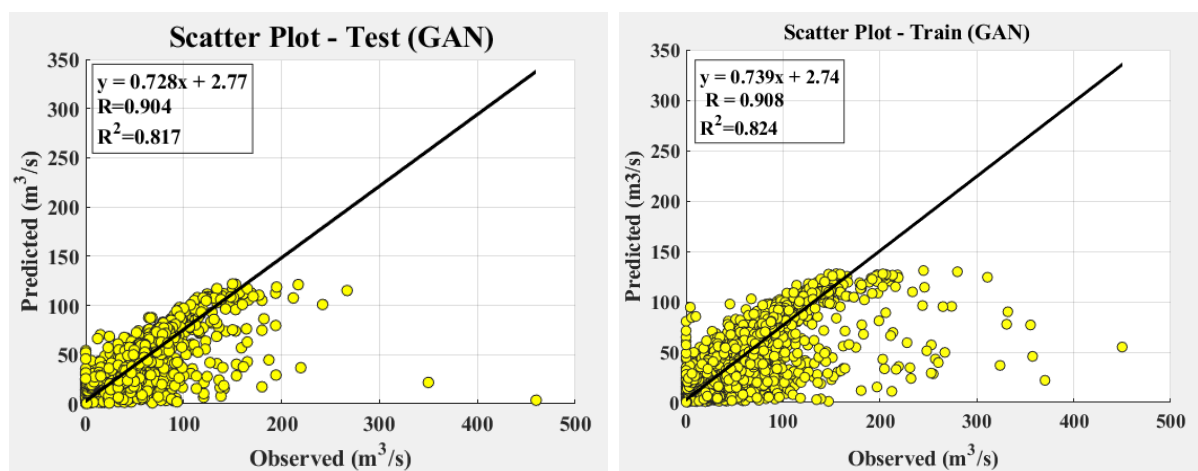
شکل ۷. نمودار خطای پیش‌بینی مدل SVR

Fig. 7 Prediction error plot of SVR model



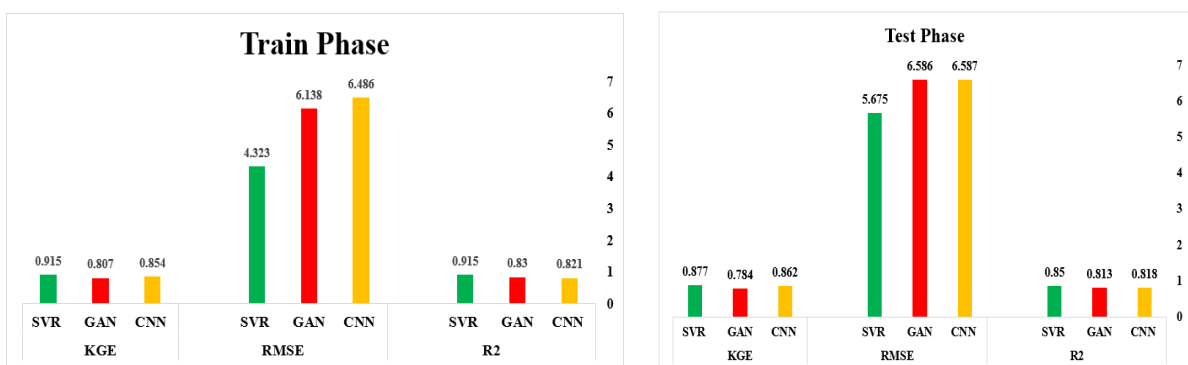
شکل ۸. نمودار پراکندگی بین داده‌های مشاهداتی و پیش‌بینی مدل CNN در دو فاز آموزش و آزمون

Fig. 8 Scatter plot between observed and predicted values of CNN model during the training and testing phases



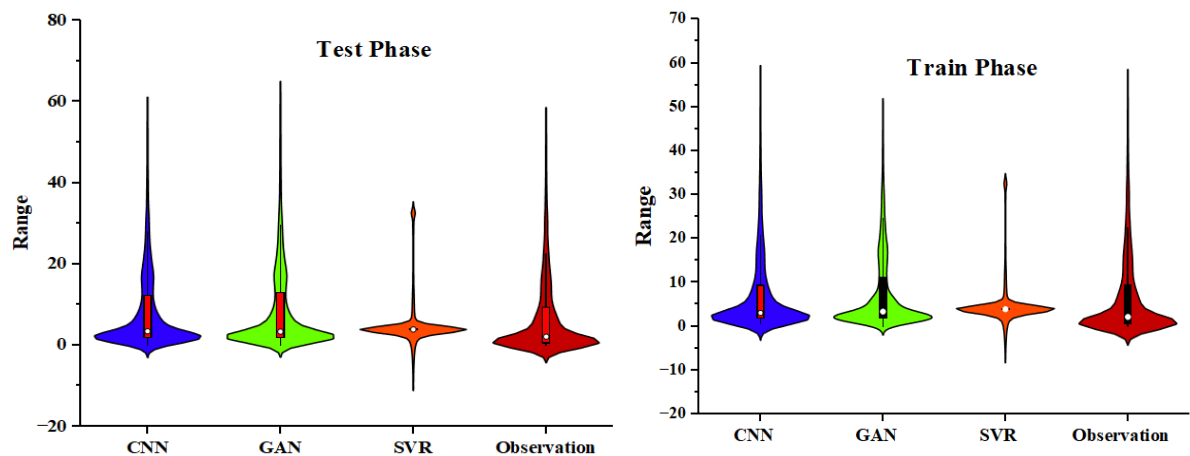
شکل ۹. نمودار پراکندگی بین داده‌های پیش‌بینی و مشاهداتی مدل GAN

Fig. 9 Scatter plot of predicted and observed values of the GAN model



شکل ۱۰. مقایسه مدل‌های پژوهش بر اساس شاخص‌های مختلف ارزیابی

Fig. 10 Comparison of the research models based on various evaluation metrics



شکل ۱۱. نمودار ویولن پلات و جعبه مدل‌های پژوهش در دو فاز آموزش و آزمون

Fig. 11 Violin and boxplot of the research models during the training and testing phases

۴- نتیجه‌گیری

این پژوهش نشان داد که دقت پیش‌بینی دبی رودخانه در حوضه‌های با نوسانات شدید هیدرولوژیک، بیش از آنکه وابسته به پیچیدگی الگوریتمی مدل باشد، به ساختار هوشمند ورودی‌ها و تطبیق مدل با ماهیت داده‌ها بستگی دارد. مدل SVR در سناریوی پنجم (شامل ترکیب بهینه‌ای از دبی‌های تأخیری و متغیرهای هواشناسی) به‌عنوان کارآمدترین مدل شناسایی شد و عملکردی برتر از CNN و GAN از خود نشان داد. برتری SVR نه تنها در کاهش خطای (RMSE (5.675 m³/s) و افزایش KGE (۰/۸۷۷)، بلکه در نزدیکی بسیار بالای PBIAS به صفر (۰/۴۷۵ درصد) آشکار شد که گواهی بر عدم سوگرایی سیستماتیک و قابلیت اطمینان آن در پیش‌بینی جریان‌های کم و زیاد است. یافته‌های این مطالعه به‌وضوح تأیید می‌کنند که حذف اطلاعات ذاتی سری زمانی جریان (دبی‌های تأخیری) حتی با استفاده از مدل‌های پیشرفته‌ی یادگیری عمیق، منجر به شکست کامل در پیش‌بینی می‌شود ($R^2 < 0.04$). این امر برجسته می‌سازد که در مدل‌سازی هیدرولوژیک داده‌محور، درک فرایندهای فیزیکی (مانند پایداری روند جریان و پاسخ تأخیری حوضه به بارش) باید هدایت‌کننده‌ی طراحی ورودی باشد، نه صرفاً دنباله‌روی از روندهای الگوریتمی. بر این اساس، پیشنهاد می‌گردد در حوضه‌های با داده‌های محدود و ناپایدار، اولویت با مدلی باشد که تعادل مناسبی بین پیچیدگی و تعمیم‌پذیری برقرار می‌کند (مانند SVR یا مدل‌های ترکیبی سبک) تا از بیش‌برازش جلوگیری شود. همچنین برای افزایش قابلیت اطمینان در پیش‌بینی رویدادهای اوج، ادغام اطلاعات غیرخطی حاصل از تجزیه‌ی سیگنال (مانند EWT یا VMD) با SVR می‌تواند مسیری امیدبخش باشد. همچنین آینده‌نگری اقلیمی هیدرولوژیک مستلزم گنجاندن خروجی‌های مدل‌های اقلیمی (GCM/RCM) به‌عنوان پیش‌بینی‌کننده‌های آینده‌نگر در چارچوب‌های یادگیری ماشین است. در نهایت، تقویت شبکه‌های هیدرومتری و کاهش شکاف‌های داده‌ای به‌ویژه در ایستگاه‌های کم‌داده‌ی ایران، باید در اولویت‌های سیاستی قرار گیرد، چرا که کیفیت و پیوستگی داده، پایه‌ی اساسی هر مدل پیش‌بینی‌کننده‌ی معتبر است. همچنین با توجه به قابلیت تعمیم‌پذیری مدل SVR در سایر ایستگاه‌های مشابه و سادگی نسبی پیاده‌سازی آن (در مقایسه با مدل‌های عمیق)، پیشنهاد می‌شود یک پلتفرم مبتنی هوشمند برخط توسعه یابد که خروجی‌های پیش‌بینی را به‌صورت شفاف و بلادرنگ در اختیار ذینفعان از جمله دهیاری‌ها، کشاورزان، سازمان آب منطقه‌ای و نهادهای نظارتی قرار دهد. چنین پلتفرمی نه تنها شفافیت تصمیم‌گیری را افزایش می‌دهد، بلکه با آگاهی‌بخشی به کشاورزان درباره‌ی احتمال نوسانات جریان، کاهش تنش آبی و بهینه‌سازی الگوی کشت را تسهیل نموده و گامی در جهت حکمرانی آب مشارکتی و مقاوم در برابر تغییرات اقلیمی ایجاد خواهد نمود. در مجموع، این پژوهش گواهی بر این اصل است که در حکمرانی آب در بحران، هوش مصنوعی نباید جانشین درک هیدرولوژیک شود، بلکه باید به‌عنوان ابزاری برای تقویت آن

مورد استفاده قرار گیرد، همان‌گونه که در این مطالعه، ترکیب هوشمندانه‌ی فیزیک حوضه و الگوریتم مناسب، راه‌حلی قابل اعتماد برای یکی از چالش‌های کلیدی مدیریت منابع آب در شمال غرب ایران ارائه کرد.

۵- سپاس‌گزاری

نویسندگان این پژوهش مراتب تشکر و قدردانی خود را از شرکت مدیریت منابع آب ایران بابت اشتراک‌گذاری اطلاعات اعلان می‌دارند. همچنین از نشریه اقلیم و بوم‌سازگان مناطق خشک و نیمه‌خشک نیز بابت تسریع در فرایند داوری و پاسخگویی بهنگام تشکر و قدردانی می‌نمایند.

۶- داده‌ها و اطلاعات

مبنای داده‌ها و اطلاعات مقاله حاضر، نویسنده اول و مسئول مکاتبات است.

۷- تعارض منافع

در این مقاله، تعارض منافی وجود ندارد و این مسأله مورد تأیید همه نویسندگان است.

۸- مشارکت نویسندگان

مشارکت نویسندگان در این مقاله به شرح زیر است:

مشارکت نویسنده اول در ایده‌پردازی پژوهش، طراحی چارچوب تحقیق، جمع‌آوری داده‌ها، تحلیل و تفسیر نتایج، و نگارش پیش‌نویس اولیه مقاله می‌باشد. مشارکت نویسنده دوم در طراحی روش تحقیق، تحلیل آماری داده‌ها، اعتبارسنجی نتایج و همکاری در نگارش و بازبینی علمی مقاله است. مشارکت نویسنده سوم در نظارت و راهنمایی بر روند انجام پژوهش، بررسی و کنترل صحت نتایج، و ویرایش نهایی متن مقاله است و در نهایت مشارکت نویسنده چهارم، تأمین داده‌ها و امکانات پژوهشی، پشتیبانی فنی و اجرایی، و مشارکت در تفسیر نتایج و بهبود ساختار مقاله می‌باشد.

۹- اصول اخلاقی

نویسندگان، اصول اخلاقی را در انجام و انتشار این اثر علمی رعایت نموده‌اند و این موضوع مورد تأیید همه آنها می‌باشد.

۱۰- حمایت مالی

این مقاله حاصل از یک پژوهش آزاد می‌باشد که هیچ گونه وابستگی مالی به دانشگاه یا سازمان و دستگاه‌های اجرایی و مؤسسات آموزش عالی نداشته و از لحاظ مالی مستقل می‌باشد.

۱۱- مراجع

1. Adnan, R. M., Mostafa, R. R., Elbeltagi, A., Yaseen, Z. M., Shahid, S., & Kisi, O. (2022). Development of new machine learning model for streamflow prediction: Case studies in Pakistan. *Stochastic Environmental Research and Risk Assessment*, 36(4), 999–1033.
2. Avand, M., Moradi, H. R., & Hazbavi, Z. (2024). Interactive changes in climatic and hydrological droughts, water quality, and land use/cover of Tajan watershed, Northern Iran. *Water*, 16(13), 1784.
3. Ayana, Ö., Kanbak, D. F., Kaya Keleş, M., & Turhan, E. (2023). Monthly streamflow prediction and performance comparison of machine learning and deep learning methods. *Acta Geophysica*, 71(6), 2905–2922.
4. 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.
5. Cho, K., & Kim, Y. (2022). Improving streamflow prediction in the WRF-Hydro model with LSTM networks. *Journal of Hydrology*, 605, 127297.

6. Cortes, C., & Vapnik, V. (1995). Support-vector networks. *Machine Learning*, 20(3), 273–297.
7. Danesh, M., Gharehbaghi, A., Mehdizadeh, S., & Danesh, A. (2025). A comparative assessment of machine learning and deep learning models for the daily river streamflow forecasting. *Water Resources Management*, 39(4), 1911–1930.
8. Ghimire, S., Yaseen, Z. M., Farooque, A. A., Deo, R. C., Zhang, J., & Tao, X. (2021). Streamflow prediction using an integrated methodology based on convolutional neural network and long short-term memory networks. *Scientific Reports*, 11(1), 17497.
9. Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., & Bengio, Y. (2020). Generative adversarial networks. *Communications of the ACM*, 63(11), 139–144.
10. Grey, V., Fletcher, T., Smith-Miles, K., Hatt, B., & Coleman, R. (2025). Harnessing the strengths of machine learning and geostatistics to improve streamflow prediction in ungauged basins; the best of both worlds. *Journal of Hydrology*, 133936.
11. Ikram, R. M. A., Hazarika, B. B., Gupta, D., Heddam, S., & Kisi, O. (2023). Streamflow prediction in mountainous region using new machine learning and data preprocessing methods: a case study. *Neural Computing and Applications*, 35(12), 9053–9070.
12. Ismail-Fawaz, A., Devanne, M., Weber, J., & Forestier, G. (2022). Deep learning for time series classification using new hand-crafted convolution filters. *2022 IEEE International Conference on Big Data (Big Data)*, 972–981.
13. Kedam, N., Tiwari, D. K., Kumar, V., Khedher, K. M., & Salem, M. A. (2024). River stream flow prediction through advanced machine learning models for enhanced accuracy. *Results in Engineering*, 22, 102215.
14. Khosravi, K., Golkarian, A., & Tiefenbacher, J. P. (2022). Using optimized deep learning to predict daily streamflow: A comparison to common machine learning algorithms. *Water Resources Management*, 36(2), 699–716.
15. Krichen, M. (2023). Generative adversarial networks. *2023 14th International Conference on Computing Communication and Networking Technologies (ICCCNT)*, 1–7.
16. Kumar, V., Kedam, N., Sharma, K. V., Mehta, D. J., & Caloiero, T. (2023). Advanced machine learning techniques to improve hydrological prediction: A comparative analysis of streamflow prediction models. *Water*, 15(14), 2572.
17. Kumshe, U. M. M., Abdulhamid, Z. M., Mala, B. A., Muazu, T., Muhammad, A. U., Sangary, O., Ba, A. F., Tijjani, S., Adam, J. M., & Ali, M. A. H. (2024). Improving short-term daily streamflow forecasting using an autoencoder based CNN-LSTM model. *Water Resources Management*, 38(15), 5973–5989.
18. Latif, S. D., & Ahmed, A. N. (2023). Streamflow prediction utilizing deep learning and machine learning algorithms for sustainable water supply management. *Water Resources Management*, 37(8), 3227–3241.
19. López-Chacón, S. R., Salazar, F., & Bladé, E. (2025). Hybrid physically based and machine learning model to enhance high streamflow prediction. *Hydrological Sciences Journal*, 70(2), 311–333.
20. Merufinia, E., Sharafati, A., Abghari, H., & Hassanzadeh, Y. (2023). On the simulation of streamflow using hybrid tree-based machine learning models: A case study of Kurkursar basin, Iran. *Arabian Journal of Geosciences*, 16(1), 28.
21. Meshram, S. G., Meshram, C., Santos, C. A. G., Benzougagh, B., & Khedher, K. M. (2022). Streamflow prediction based on artificial intelligence techniques. *Iranian Journal of Science and Technology, Transactions of Civil Engineering*, 46(3), 2393–2403.
22. Miao, S., & Hung, W. H. (2020). River flooding forecasting and anomaly detection based on deep learning. *Ieee Access*, 8, 198384–198402.
23. Naabil, E., Lampitey, B. L., Arnault, J., Olufayo, A., & Kunstmann, H. (2017). Water resources management using the WRF-Hydro modelling system: Case-study of the Tono dam in West Africa. *Journal of Hydrology: Regional Studies*, 12, 196–209.
24. Nasir, N., Irwan, D., Ahmed, A. N., Ibrahim, S. L., Ibrahim, I., Sherif, M., & El-Shafie, A. (2025). Harnessing machine learning for streamflow prediction: A comparative study of advanced models in the Upper Klang River Basin, Malaysia. *Journal of Hydrology: Regional Studies*, 60, 102565.
25. Saeidi, P., Mehrdadi, N., Karbassi, A., & Ardestani, M. (2022). Integrated river-estuary modeling to assess spawning and habitation area for Caspian WhiteFish in response to upstream pollution

- sources (case study: Tajan River estuary). *International Journal of Environmental Research*, 16(5), 93.
26. Solanki, H., Vegad, U., Kushwaha, A., & Mishra, V. (2025). Improving streamflow prediction using multiple hydrological models and machine learning methods. *Water Resources Research*, 61(1), e2024WR038192.
 27. Sun, K., Rajabtabar, M., Samadi, S., Rezaie-Balf, M., Ghaemi, A., Band, S. S., & Mosavi, A. (2021). An integrated machine learning, noise suppression, and population-based algorithm to improve total dissolved solids prediction. *Engineering Applications of Computational Fluid Mechanics*, 15(1), 251–271.
 28. 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.
 29. Vinokić, L., Dotlić, M., Prodanović, V., Kolaković, S., Simonovic, S. P., & Stojković, M. (2025). Effectiveness of three machine learning models for prediction of daily streamflow and uncertainty assessment. *Water Research X*, 27, 100297.
 30. Wang, L., Guo, Y., & Fan, M. (2022). Improving annual streamflow prediction by extracting information from high-frequency components of streamflow. *Water Resources Management*, 36(12), 4535–4555.
 31. 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.
 32. Xu, W., Chen, J., Zhang, X. J., Xiong, L., & Chen, H. (2022). A framework of integrating heterogeneous data sources for monthly streamflow prediction using a state-of-the-art deep learning model. *Journal of Hydrology*, 614, 128599.
 33. Yu, X., Wang, Y., Wu, L., Chen, G., Wang, L., & Qin, H. (2020). Comparison of support vector regression and extreme gradient boosting for decomposition-based data-driven 10-day streamflow forecasting. *Journal of Hydrology*, 582, 124293.
 34. Zhou, Y., Duan, Y., Yao, H., Li, X., & Li, S. (2025). Incorporating hydrological constraints with deep learning for streamflow prediction. *Expert Systems with Applications*, 259, 125379.