



Estimating the Suspended Sediment Load in the Sefidrood River Using Meta-Heuristic Models

Milad Fardadi Shilsar¹, Ebrahim Amiri^{2*}, Jalal Behzadi³

1-PhD Student, Department of Water Engineering, La, C., Islamic Azad University, Lahijan, Iran.

2-Professor, Department of Water Engineering, La, C., Islamic Azad University, Lahijan, Iran.

3-Assistant Professor, Department of Agriculture, La, C., Islamic Azad University, Lahijan, Iran.

Received:2026-06-10

Accepted:2026-09-19

Revised:2026-09-16

Published online:2026-06-20

*Corresponding Author's Email: o.bazrafsan@hormozgan.ac.ir

ARTICLE INFO

ABSTRACT

Keywords:

Suspended Sediment Flow Estimation, Artificial Neural Network, Modeling, Metaheuristic Models, Machine Learning Models,

Introduction

Suspended sediments are fine particles transported within the water column under turbulent flow conditions, and their movement in rivers is controlled by complex interactions between hydrological processes, flow characteristics, and sediment properties. The deposition of these sediments can significantly reduce river conveyance capacity, increasing flood vulnerability and affecting sustainable water resource management. Rivers, as vital natural systems, are continuously influenced by the transport and accumulation of sediments. Historically, suspended sediment estimation has relied on empirical, statistical, and physically based models. However, in recent years, data-driven approaches, particularly Artificial Neural Networks (ANN) and evolutionary optimization models, have demonstrated considerable potential for improving prediction accuracy. Hybrid ANN models, including ANN-Genetic Algorithm (GA), ANN-Particle Swarm Optimization (PSO), and ANN-Imperialist Competitive Algorithm (ICA), have been widely applied to enhance model optimization, improve convergence performance, and increase the capability of nonlinear prediction. Furthermore, nonlinear machine learning techniques, such as Random Forest (RF) and Support Vector Machine (SVM), have also shown promising performance in suspended sediment load estimation compared with conventional approaches. This research evaluates a comprehensive set of six models—ANN, ANN-GA, ANN-PSO, ANN-ICA, SVM, and RF—for estimating suspended sediment load in the Sefidrood River. The Sefidrood River is of great importance in northern Iran due to its high sediment transport capacity, considerable rainfall regime, and key role in regional water supply and sediment transfer to the Caspian Sea. The main objective is to compare the predictive performance and generalization capability of different data-driven models to support effective sediment management and future water resource decisions in this region.

Materials and Methods

The Sefidrood River, the largest river in Northern Iran, was selected as the study area. Originating from the Zanzan and Shahrood mountains, the river extends over 765 km, draining a vast catchment area of approximately 59,400 km² into the Caspian Sea. It is characterized by high discharge variability, considerable sediment transport capacity, and a high potential for flooding and sedimentation.

For the simulation, daily discharge and suspended sediment load data were collected from the Gilan Regional Water Authority, covering the statistical period from 2007 to 2022 (1386–1401). The dataset was

How to cite:

Fardadi Shilsar, M. Amiri, E. Behzadi, J. *Estimating the Suspended Sediment Load in the Sefidrood River Using Meta-Heuristic Models*. Journal of Hydraulics and Water Science, 36 (1):65-82

<https://doi.org/10.22034/hws.2026.72908.1057>



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



randomly partitioned, with 70% used for model training and the remaining 30% reserved for independent testing. Prior to modeling, statistical outliers were treated, and the data were normalized using Equation (1) to improve data quality and enhance simulation performance. Six data-driven models were implemented to estimate suspended sediment load. These included the standard ANN model, three hybrid meta-heuristic versions (ANN-GA, ANN-PSO, and ANN-ICA), and two machine learning techniques (SVM and RF). The models were trained using the available dataset, and their performance was evaluated based on prediction accuracy criteria, including the Root Mean Square Error (RMSE), Mean Squared Error (MSE), and coefficient of determination (R^2). The base ANN model was configured with one hidden layer containing 10 neurons, employing the 'tansig' and 'purelin' activation functions for the hidden and output layers, respectively, over 1000 epochs. The hybrid optimizers were set with specific parameters; GA used a population size of 50 with 200 iterations, PSO was implemented with 50 particles and 200 iterations, and ICA utilized 100 countries and 30 imperialists for 150 iterations. Finally, the models were compared using both training and independent testing results to assess predictive accuracy and generalization capability.

Results and Discussion

The performance of the six data-driven models was evaluated using the coefficient of determination (R^2), Mean Squared Error (MSE), and Root Mean Square Error (RMSE). During the training phase, the Random Forest (RF) model achieved the highest predictive accuracy ($R^2 = 0.90$ and $RMSE = 3.49$), followed by ANN-GA ($R^2 = 0.87$), SVM ($R^2 = 0.86$), ANN ($R^2 = 0.85$), ANN-ICA ($R^2 = 0.84$), and ANN-PSO ($R^2 = 0.76$). In the independent testing phase, ANN-ICA produced the best overall performance with the lowest RMSE (4.14) and an R^2 value of 0.86. ANN-GA achieved the same R^2 (0.86) with a slightly higher RMSE (4.61), while SVM also demonstrated satisfactory predictive capability ($R^2 = 0.85$ and $RMSE = 4.37$). Although RF showed excellent performance during training, its testing accuracy decreased ($R^2 = 0.83$), indicating relatively lower generalization capability than ANN-ICA. Overall, the obtained R^2 values ranging from approximately 0.76 to 0.90 indicate that all developed models were capable of capturing the nonlinear relationship between discharge and suspended sediment load. To complement the conventional performance indices, the predictive errors of the six models were statistically compared using the non-parametric Friedman test. The test revealed a highly significant difference among model performances ($p = 2.0 \times 10^{-8} < 0.05$), indicating that at least one model differed significantly from the others. The mean-rank analysis showed that ANN obtained the lowest average error rank (2.04), followed by ANN-PSO (3.14), ANN-GA (3.86), ANN-ICA (3.88), SVM (3.88), and RF (4.20). This ranking differs from the RMSE- and R^2 -based evaluation because the Friedman test considers the ranking of prediction errors at every observation rather than relying on a single summary statistic. Therefore, a model with a lower overall RMSE may still receive a higher mean rank if it produces relatively larger errors at some time steps. Pairwise comparisons were subsequently performed using the Bonferroni post-hoc procedure. The results indicated statistically significant differences between the conventional ANN model and ANN-GA, ANN-ICA, RF, and SVM ($p < 0.001$), whereas the comparison between ANN and ANN-PSO was marginally non-significant ($p = 0.053$). No statistically significant differences were detected among ANN-GA, ANN-PSO, ANN-ICA, RF, and SVM ($p > 0.05$), suggesting that these advanced data-driven approaches provide statistically comparable predictive performance despite slight numerical differences in RMSE and R^2 .

Conclusion

This study evaluated six data-driven approaches, including ANN, ANN-GA, ANN-PSO, ANN-ICA, RF, and SVM, for suspended sediment load estimation in the Sefidrood River. The results demonstrated that model performance differed between the training and testing stages. RF achieved the highest predictive accuracy during training, whereas ANN-ICA provided the best generalization performance during testing by producing the lowest RMSE together with a high coefficient of determination. ANN-GA and SVM also achieved competitive prediction accuracy, confirming the effectiveness of advanced machine learning and hybrid optimization techniques for nonlinear suspended sediment modeling. Beyond the conventional evaluation criteria, the Friedman statistical test confirmed that the overall differences among the six models were statistically significant ($p < 0.05$). However, the Bonferroni post-hoc analysis showed that these differences were mainly associated with the conventional ANN model, while no significant differences were observed among ANN-GA, ANN-PSO, ANN-ICA, RF, and SVM. These findings demonstrate that model selection should not rely solely on RMSE or R^2 but should also consider statistical significance tests to provide a more comprehensive assessment of predictive performance. The proposed models can serve as reliable tools for suspended sediment estimation and support sediment management and water resources planning in northern Iran. Nevertheless, the present study was conducted using daily discharge and suspended sediment data from a single hydrometric station during the 2007–2022 period. Therefore, extending the conclusions to other watersheds should be undertaken cautiously. Future research should investigate the applicability of the best-performing models in other river basins and incorporate additional explanatory variables such as rainfall, land use, and morphological characteristics, as well as advanced deep learning approaches, to further improve prediction accuracy and model generalization.



برآورد میزان بار معلق رسوبات در رودخانه سفیدرود با استفاده از مدل‌های فراابتکاری

میلاد فردادی شیل سر^۱، ابراهیم امیری^{۲*}، جلال بهزادی^۳

۱- دانشجوی دکتری، گروه مهندسی آب، واحد لاهیجان، دانشگاه آزاد اسلامی، لاهیجان، ایران.

۲- استاد، گروه مهندسی آب، واحد لاهیجان، دانشگاه آزاد اسلامی، لاهیجان، ایران.

۳- استادیار، گروه کشاورزی، واحد لاهیجان، دانشگاه آزاد اسلامی، لاهیجان، ایران.

تاریخ پذیرش: ۱۴۰۵/۰۶/۲۸

تاریخ دریافت: ۱۴۰۵/۰۳/۲۰

تاریخ انتشار آنلاین: ۱۴۰۵/۰۳/۳۱

تاریخ ویرایش: ۱۴۰۵/۰۶/۲۵

E-mail: o.bazrafsan@hormozgan.ac.ir * نویسنده مسئول

چکیده

این پژوهش کاربرد تلفیق شبکه عصبی مصنوعی با الگوریتم‌های بهینه‌سازی را به منظور برآورد کمی بار معلق رسوبی ایستگاه آستانه رودخانه سفیدرود نشان داد. داده‌های ورودی این شبیه‌سازی شامل دبی جریان و بار رسوب معلق رودخانه مذکور طی دوره آماری ۱۳۸۶ تا ۱۴۰۱ بود. شش مدل شبیه‌سازی شامل شبکه عصبی مصنوعی، نسخه‌های ترکیبی آن با مدل‌های ژنتیک، ازدحام ذرات و رقابت استعماری، و همچنین دو مدل یادگیری ماشین شامل جنگل تصادفی و بردار پشتیبان به کار گرفته شد. داده‌ها به صورت تصادفی به دو بخش آموزش (۷۰ درصد) و آزمایش (۳۰ درصد) تقسیم شدند، تا کارایی و عملکرد مدل‌ها به درستی اعتبارسنجی شود. نتایج نشان داد که مدل ANN-ICA عملکرد مطلوبی را با کمترین مقدار RMSE (۴/۱۴ میلی گرم بر لیتر) و ضریب تعیین ۰/۸۶ در مرحله آزمایش، بهترین قابلیت تعمیم را در بین مدل‌های بررسی شده نشان داد. همچنین مدل ANN-GA با ضریب تعیین ۰/۸۶ عملکرد مشابهی داشت. اگرچه مدل جنگل تصادفی در مرحله آموزش بالاترین دقت را نشان داد ($R^2=0.90$) و RMSE برابر با ۳/۴۹ میلی گرم بر لیتر، اما کاهش عملکرد آن در مرحله آزمایش ($R^2=0.83$ و RMSE برابر با ۴/۶ میلی گرم بر لیتر) احتمال بیش‌برازش را نشان داد. همچنین، نتایج آزمون ناپارامتری فریدمن وجود اختلاف آماری معنادار بین عملکرد مدل‌ها را تأیید کرد ($p<0.05$) و آزمون تعقیبی بونفرونی نشان داد که اختلاف معنادار عمدتاً بین مدل پایه ANN و سایر مدل‌های پیشرفته مشاهده شد، در حالی که بین مدل‌های پیشرفته اختلاف آماری معناداری وجود نداشت. تحلیل الگوی خطا و آزمایش‌های آماری نیز بیانگر نبود بایاس^۱ زمانی و فقدان خطای ساختاری در شبیه‌سازی بود.

کلمات کلیدی: برآورد دبی معلق رسوب، شبکه عصبی مصنوعی، مدل‌سازی، مدل‌های فراابتکاری، مدل‌های یادگیری ماشین.

^۱ بایاس (Bias) یا سوگیری به هر نوع خطای سیستماتیک گفته می‌شود.

۱- مقدمه

رسوبات معلق به ذرات ریز منتقل شده توسط جریان آب اطلاق می‌شوند که آن‌قدر کوچک‌اند که در اثر آشفتگی جریان، به‌جای ته‌نشینی در بستر، در ستون آب رودخانه باقی می‌مانند (Parsons et al., 2022; Dethier et al., 2022; Hoffmann et al., 2023; al., 2015). این ذرات می‌توانند ماهیت فیزیکی داشته باشند؛ مانند ذرات ریز کانی، مواد آلی یا میکروپلاستیک‌ها (Emami et al., 2024; Saadat et al., 2025a; Bezak et al., 2025). فرآیند عناصر سمی باشند (Fawell & Nieuwenhuijsen, 2003). فرآیند انتقال این رسوبات در رودخانه‌ها براساس معادله حاکم انتقال جرم (املاح) و حل عددی یا تحلیلی آن توصیف می‌شود (Fardadi Shilsar et al., 2023; Gulliver, 2007; Saadat et al., 2022; Zangeneh Asadi et al., 2025). در این معادلات، پدیده جابجایی مهم‌ترین مکانیسم انتقال به‌شمار می‌رود، زیرا می‌تواند موجب حرکت پایدار رسوبات به سمت پایین‌دست، انباشت موضعی آن‌ها و تغییر تراز بستر شود (Fardadi Shilsar et al., 2022c; Saadat & Mazaheri, 2024b; Zhou et al., 2023). که در بلندمدت با کاهش ظرفیت مقطع و افزایش سطح آب، منجر به تشدید دبی اوج سیلاب و مخاطرات آن خواهد شد (Saadat et al., 2024; Vázquez-Tarrío et al., 2024). از مهم‌ترین پیامدهای ته‌نشینی رسوبات رودخانه‌ای می‌توان به کاهش ظرفیت عبور سیلاب، ناپایداری مورفولوژیکی مجرا و افزایش خسارت‌های ناشی از رخدادهای حدی اشاره کرد (Vanoni, 2014; Knighton, 2006). از این رو، رودخانه‌ها را می‌توان یکی از ارزشمندترین و در عین حال آسیب‌پذیرترین سامانه‌های طبیعی دانست که هرگز به‌طور کامل عاری از آلودگی نبوده و همواره تحت تأثیر ورود املاح و رسوبات ناشی از فعالیت‌های انسانی و تغییر کاربری اراضی هستند (Smits et al., 2000; Fardadi Shilsar et al., 2022a, 2022b; Saadat & Mazaheri, 2024a; Saadat et al., 2025b; Dethier et al., 2022).

روش‌های مرسوم در برآورد رسوب شامل مدل‌های تجربی (مانند منحنی سنج رسوب)، مدل‌های آماری، مدل‌های فیزیکی و مدل‌های داده‌محور از جمله شبکه‌های عصبی مصنوعی (ANN) و مدل‌های بهینه‌سازی تکاملی هستند. در دهه گذشته، شبکه‌های عصبی و یادگیری ماشین به عنوان ابزارهای اصلی برای افزایش دقت برآورد رسوب شناخته شده‌اند و عملکرد بهتری نسبت به مدل‌های کلاسیک دارند (Nda et al., 2023; Aires et al., 2023; Joshi et al., 2024).

مدل شبکه عصبی مصنوعی (ANN) تاکنون در مطالعات بسیاری در مبحث رسوبات رودخانه‌ای به کار گرفته شده است. پارامترهای مهم آن شامل تعداد لایه‌های نورون، انتخاب ورودی‌های بهینه، نوع تابع انتقال، مدل آموزش، استفاده از مدل‌های بهینه‌سازی، و پیش‌پردازش

داده‌ها است. لازم به ذکر است که انتخاب ورودی مناسب و جلوگیری از بیش‌برازش در آن حائز اهمیت است (Hasanpour Kashani et al., 2014; Kisi & Shiri, 2012). برای اولین بار از شبکه عصبی ANN در برآورد رسوب معلق رودخانه‌ای استفاده کردند. در دو دهه مدل‌سازی رسوب با هوش مصنوعی، (Allawi et al., 2023) بارها شبکه عصبی ANN را به عنوان دقیق‌ترین ابزار بازسازی داده‌های رسوب معلق در شرایط مختلف آب‌وهوایی معرفی کرده‌اند. (Le et al., 2024) نشان دادند که استفاده از شبکه‌های یادگیری ماشین مانند ANN برای کاهش عدم قطعیت در پیش‌بینی رسوب بستر رودخانه منجر به نتایج دقیق‌تری شده است. در مطالعات (Le et al., 2024) روی پیش‌بینی بار بستر رودخانه، و (Emami & Choopan, 2019) روی تخمین بار معلق رسوب دریاچه ارومیه، مدل‌های ترکیبی ANN-PSO، ANN-GA و ANN-ICA توانستند در داده‌های همان سامانه‌های رودخانه‌ای (داده‌های آزمایشگاهی و میدانی هر مطالعه) میانگین خطای RMSE را به حدود ۳ تا ۵ درصد کاهش دهند و در مقایسه با مدل‌های منفرد، همگرایی سریع‌تر، تعمیم‌پذیری بهتر و مقاومت بیشتری در برابر مینیمم موضعی نشان دهند. (Naseri et al., 2019) مدل بهینه‌سازی شده GA را برای رودخانه‌های گرگانود (ایستگاه قزاقلی) و فریمان (ایستگاه باغ‌عباسی) با دقت بالاتر نسبت به روش منحنی سنج رسوب (^{1}SRC) گزارش کردند. برخی پژوهشگران ترکیب تابع پایه شعاعی (^{2}RBF) و الگوریتم ژنتیک (GA) را برای بهبود انتخاب ورودی و افزایش دقت پیش‌بینی ماهانه رسوب به کار گرفتند (Teixeira et al., 2020). (Yadav et al., 2022) در تحقیقی مدل ANN-GA را برای برآورد رسوبات معلق رودخانه تیکاراپارا در هند به کار بردند و به نتیجه برتری مدل ترکیبی نسبت به مدل ANN و SRC رسیدند. در زمینه مدل‌های ANN-ICA، (Nikoo et al., 2018) نیز تلفیق ICA و ANN را برای پیش‌بینی رسوبات معلق رودخانه کارون با افزایش انعطاف‌پذیری و دقت مدل تأیید کردند. علاوه بر آن، در تحقیق دیگر نیز، تأثیر و کاربرد مدل ANN-ICA را در برآورد رسوب رودخانه‌ای با کاهش هزینه محاسباتی و دقت بالا گزارش کردند (Khanduzi et al., 2024). (Yadav et al., 2018) نشان دادند که مدل ANN-GA برای برآورد رسوب معلق رودخانه ماهانادی هند، علاوه بر دقت بالا، به بهینه‌سازی همزمان پارامترها منجر می‌شود. این در حالی است که (Collins et al., 2010) از مدل GA برای اولین بار در مدل‌سازی رسوب استفاده کردند و پس از آن مدل‌های ترکیبی متعددی برای برآورد دقیق و کاهش خطا ارائه کردند.

مدل‌های جنگل تصادفی (RF) و ماشین بردار پشتیبان (SVM) به عنوان نسل جدید مدل‌های داده‌محور غیرخطی، جایگاه ویژه‌ای در مدل‌سازی بار معلق رسوبی پیدا کرده‌اند. در حوزه رسوب‌گذاری رودخانه‌ای، نخستین استفاده جدی از SVM برای پیش‌بینی بار معلق

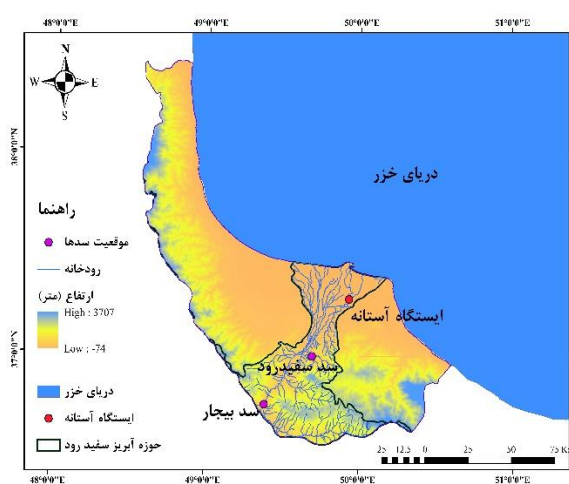
² Radial Basis Function¹ Sediment Rating Curve

رودخانه سفیدرود ارزیابی شده‌اند و علاوه بر دقت، تعمیم‌پذیری، رفتار خطا و پایداری زمانی مدل‌ها نیز به‌طور نظام‌مند بررسی شده است؛ بدین ترتیب، یک چارچوب عملی برای انتخاب مدل‌های رسوب‌محور در مدیریت حوضه‌های پررسوب شمال ایران ارائه می‌شود. علاوه بر این، رودخانه سفیدرود به‌دلیل ویژگی‌های هیدرولوژیکی، بارندگی بالا، رسوب‌خیزی چشمگیر و نقش کلیدی در تأمین آب استان و انتقال رسوب به دریای خزر اهمیت ویژه‌ای دارد. با توجه به تنوع اقلیم و شرایط فیزیوگرافی رودخانه‌های گیلان، اعتبارسنجی مدل‌ها برای رودخانه فوق‌الذکر ضروری است. به‌طور کلی، استفاده از مدل‌های داده‌محور می‌تواند به کاهش عدم قطعیت در برآورد رسوب در محیط‌های شمال ایران کمک کند. با وجود این، باید توجه داشت که نتایج این پژوهش براساس داده‌های روزانه یک ایستگاه منتخب در یک بازه زمانی تقریباً ۱۵ ساله به دست آمده و تعمیم‌پذیری آن به سایر رودخانه‌ها و بازه‌های زمانی نیازمند بررسی‌های تکمیلی است که این موضوع یکی از محدودیت‌های مهم این مطالعه است.

۲- مواد و روش‌ها

۲-۱- معرفی محدوده پژوهش

رودخانه سفیدرود بزرگ‌ترین رود شمال ایران با سرچشمه از کوه‌های زنجان و شاهرود، طول ۷۶۵ کیلومتر، و ارتفاع مصب حدود منفی ۲۸ متر است (شکل ۱). حوضه آبریز سفیدرود با وسعت حدود ۵۹,۴۰۰ کیلومتر مربع، جزو زیرحوضه‌های دریای خزر است و عمده رودخانه‌های گیلان را به خود متصل می‌کند. رودخانه فوق‌الذکر، دائمی، پرآب، با شیب زیاد، بارش بالا، مستعد سیل و رسوب‌دهی فراوان، و از مهم‌ترین منبع آب و رسوب شمال ایران به شمار می‌رود.



شکل ۱. محل قرارگیری رودخانه سفیدرود در استان گیلان

Fig 1. Location of the Sefidrood River in Gilan Province

۲-۲- جمع‌آوری داده‌ها

داده‌های مورد استفاده مربوط به دوره آماری ۱۳۸۶ تا ۱۴۰۱ در ایستگاه هیدرومتری آستانه بر روی رودخانه سفیدرود در استان

در حوضه کاکارضا لرستان گزارش شده است که نشان داد این مدل نسبت به منحنی سنجه کلاسیک دقت به‌مراتب بالاتری دارد (Shahi & Dehghani, 2017). در مورد RF نیز به‌کارگیری این مدل در گرگان‌رود بیانگر توان آن در بازتولید روابط پیچیده دبی-رسوب و کاهش محسوس خطا نسبت به روش‌های تجربی متداول بود (Seyyedi et al., 2017). مدل SVM به‌واسطه استفاده از کرنل‌های غیرخطی و حاشیه بهینه، قادر است الگوهای پیچیده و غیرخطی موجود در داده‌های رسوبی را با حداقل تعداد پارامتر تنظیمی شناسایی کند. نتایج پژوهش‌های مختلف نشان می‌دهد که این مدل در بسیاری از موارد از رگرسیون‌های کلاسیک و حتی برخی ساختارهای شبکه عصبی عملکرد بهتری داشته و مقادیر R^2 تا حدود ۰/۹۸ و RMSE بین ۰/۱۵ تا ۰/۲۴ (مقدار نرمال‌سازی شده) برای آن ثبت شده است (Moradinejad, 2024; Shahi Nejad & Dehghani, 2017). به عنوان نمونه، در مطالعه کاکارضا، SVM موفق شد بار معلق را با $R^2=0.867$ و $RMSE=0.024$ تن بر روز برآورد کند، و در ایستگاه پل دواب قره‌چای نیز کارایی این مدل با گزارش R^2 بزرگ‌تر از ۰/۹۸ در مقایسه با GMDH تأیید شد (Moradinejad, 2024). نتایج Meshram et al. (2020) نیز نشان می‌دهد که دقت SVM در تخمین غلظت رسوب، هم‌تراز مدل‌های ترکیبی پیشرفته مانند ANFIS-based سیستم‌های فازی-عصبی است. در مقابل، مدل RF با رویکرد اجتماع درختی و نمونه‌برداری تصادفی از متغیرها، چارچوبی پایدار برای تحلیل داده‌های هیدرولوژیکی پرنویز فراهم می‌کند و در بسیاری از مطالعات رسوبی به عنوان مدلی با حساسیت کمتر به داده‌های پرت معرفی شده است. کاربرد این مدل در حوضه گرگان‌رود نشان داد که ترکیب چندین درخت تصمیم می‌تواند RMSE را برای ایستگاه‌های با دبی بالا به مقادیر نسبتاً پایینی در حد ده‌ها تن بر روز کاهش دهد و بر منحنی سنجه سنتی برتری یابد (Seyyedi et al., 2017). همچنین، Kisi et al. (2024) در حوضه یانگ‌تسه عملکرد RF را به‌طور سیستماتیک با RT و M5P مقایسه کرده و برتری RF را از نظر شاخص‌های خطا گزارش کردند، در حالی که AIDahoul et al. (2022) نشان دادند این مدل در مقیاس ماهانه نیز برای برآورد رسوب، در کنار روش‌هایی مانند RNN و GRU، گزینه‌ای قابل اتکا محسوب می‌شود. اگر لازم شد می‌توان این بدنه متنی را با مطالعات دیگری درباره کاربرد RF و SVM در مدل‌سازی بار معلق، بسته به محدودیت تعداد رفرنس مقاله، غنی‌تر نیز کرد.

با توجه به پژوهش‌های اخیر می‌توان اذعان داشت که اغلب مطالعات در زمینه برآورد رسوبات معلق رودخانه‌ها یا بر یک کلاس مدل (مثلاً تنها ANN یا تنها SVM یا RF) تمرکز داشته‌اند و یا کمتر به ارزیابی هم‌زمان ترکیب جامع مدل‌های داده‌محور در محیط‌های پیچیده شمال ایران پرداخته‌اند. نوآوری اصلی این پژوهش در آن است که برای نخستین بار شش مدل داده‌محور شامل ANN، سه نسخه هیبرید فراابتکاری آن (ANN-ICA، ANN-PSO، ANN-GA) و دو مدل پیشرفته یادگیری ماشین (SVM و RF) به‌صورت هم‌زمان در

(et al., 2021). در پایان، برای تولید نمودارها، مقادیر نرمال‌شده (مقادیر بین صفر و یک) با اعمال رابطه معکوس نرمال‌سازی به مقادیر واقعی دبی و رسوب معلق تبدیل شدند تا نمودارهای مقایسه‌ای و تحلیل خطا براساس مقادیر فیزیکی قابل تفسیر رسم شوند.

۲-۳- معرفی مدل‌های هوش مصنوعی

در این پژوهش برای برآورد بار رسوب معلق، از شش چارچوب داده‌محور شامل شبکه عصبی مصنوعی پایه (ANN)، نسخه‌های بهینه شده آن با مدل‌های ژنتیک (ANN-GA)، ازدحام ذرات (ANN-PSO) و رقابت استعماری (ANN-ICA)، و نیز دو مدل یادگیری ماشین جنگل تصادفی (RF) و ماشین بردار پشتیبان (SVM) استفاده شد که مد نظر بسیاری از پژوهشگران بوده است (Biazar et al., 2020; Dinpashoh et al., 2017). تمامی مدل‌ها براساس کمینه‌سازی ریشه میانگین مربع خطا (RMSE) واسنجی شدند تا امکان مقایسه دقت و رفتار آن‌ها در دو حالت آموزش و آزمایش فراهم شود. به‌کارگیری RMSE به عنوان تابع هدف، رویکردی پذیرفته‌شده در ادبیات مدل‌سازی رسوب معلق است و در توسعه طیف وسیعی از مدل‌های داده‌محور و هیبرید مانند ANN، PSO-ANN، ANFIS-DE (Allawi et al., 2023; Gul et al., 2021; Syed et al., 2005; Ebtehaj & Bonakdari, 2017).

۲-۳-۱- مدل ANN

شبکه عصبی مصنوعی (ANN) در این مطالعه به‌صورت یک پرسپترون چندلایه با سه بخش اصلی ورودی، لایه پنهان و خروجی پیاده‌سازی شد که آموزش آن بر پایه روش پس‌انتشار خطا انجام گرفت. معماری نهایی پس از آزمون ساختارهای مختلف، به‌صورت یک‌لایه پنهان با ۱۰ نورون، استفاده از تابع فعال‌ساز تانژانت هائپربولیک در لایه پنهان و تابع خطی در خروجی، و ۱۰۰۰ دوره تکرار آموزش انتخاب شد که با توصیه‌های موجود در مطالعات پیشین سازگار است (Kisi & Shiri, 2012; Sanikhani et al., 2012; Razavizadeh & Dargahian, 2018). مستندسازی تغییرات RMSE در طول فرآیند یادگیری و توجیه انتخاب تعداد نورون‌ها و نوع توابع فعال‌ساز برای اطمینان از عدم بیش‌برازش انجام شد.

۲-۳-۲- مدل‌های فراابتکاری هیبرید

در مدل ترکیبی ANN-GA، مدل ژنتیک با بهینه‌سازی هم‌زمان وزن‌ها و بایاس‌ها، ساختار شبکه را به گونه‌ای تنظیم کرد که احتمال گیر افتادن در مینیمم‌های موضعی کاهش یابد و سرعت همگرایی افزایش پیدا کند (Naseri et al., 2019; Emami et al., 2018). در این پیکربندی، اندازه جمعیت ۵۰، تعداد نسل‌ها ۲۰۰ و نرخ ترکیب ۰/۸ در نظر گرفته شد، در حالی که نرخ جهش به‌صورت خودکار براساس تنظیمات پیش‌فرض تعیین شد و شبکه پایه با ۱۰ نورون در

گیلان است. این ایستگاه در مختصات جغرافیایی تقریباً ۴۹ درجه، ۵۶ دقیقه و ۲ ثانیه طول جغرافیایی و ۳۷ درجه، ۱۶ دقیقه و ۴۱ ثانیه عرض جغرافیایی با ارتفاع ۷- متر از سطح دریا واقع شده و توسط شرکت آب منطقه‌ای گیلان بهره‌برداری می‌شود. این داده‌ها شامل دبی رودخانه (مترمکعب بر روز) و بار معلق رسوبات (میلی گرم بر لیتر) است. در این دوره آماری ذکر شده، در مجموع حدود ۳۴۱ داده روزانه اولیه از دبی و بار رسوب معلق در دسترس بود (به‌طور متوسط حدود ۲۰ روز در سال). به‌منظور اطمینان از کیفیت داده‌ها، مقادیر غیرمنطقی و پرت حذف شدند و پس از بررسی آماری (بخش بعدی)، ۱۶۵ مشاهده معتبر برای تحلیل و مدل‌سازی باقی ماند. این مجموعه شامل ترکیبی از شرایط کم‌آبی، نرمال و پربابی، از جمله بخشی از رخداد‌های سیلابی دارای داده رسوب، است. لازم به ذکر است در این پژوهش، به‌دلیل عدم دسترسی به داده‌های کامل بارندگی روزانه هم‌مقیاس با ایستگاه آستانه رودخانه سفیدرود در دوره مشخص و برای حفظ سادگی و پایداری ساختار مدل، رابطه داده‌محور صرفاً میان دبی روزانه و بار رسوب معلق استخراج شد؛ با این حال، در گام‌های بعدی می‌توان متغیرهایی مانند بارندگی، کاربری اراضی و شاخص‌های مورفولوژیک را نیز برای بهبود مدل‌ها بعدها وارد کرد. برای شناسایی داده‌های پرت از یک روش دو مرحله‌ای استفاده شد. در گام اول، با استفاده از معیار فاصله چارکی (IQR^1)، نقاطی که مقدار آن‌ها در متغیرهای دبی یا رسوب معلق بزرگ‌تر از $Q_3 + 1.5 \times IQR$ یا کوچک‌تر از $Q_1 - 1.5 \times IQR$ بود، به عنوان «نامزد داده پرت» علامت‌گذاری شدند. در گام دوم، این نقاط روی نمودارهای سری زمانی و رابطه دبی-رسوب به‌صورت بصری بررسی شدند و تنها مقادیری که علاوه بر معیار آماری، از نظر هیدرولوژیکی نیز غیرمنطقی تشخیص داده شدند (مثل جهش ناگهانی رسوب بدون تغییر متناظر در دبی)، از مجموعه داده حذف شدند. سپس، داده‌های نهایی نرمال‌سازی شدند تا کیفیت داده‌ها بهبود یابد. به‌طور کلی، در شبکه‌های عصبی مصنوعی، حداقل حجم نمونه به تعداد پارامترهای قبل تنظیم (وزن‌ها و بایاس‌ها) بستگی دارد و عدد ثابتی که برای همه مسائل قابل تعمیم باشد در ادبیات گزارش نشده است. با این حال، در بسیاری از مطالعات هیدرولوژیک توصیه می‌شود تعداد داده‌های آموزش دست‌کم یک مرتبه بزرگ‌تر از تعداد وزن‌های شبکه باشد و معمولاً از چند صد تا چند هزار مشاهده روزانه برای آموزش استفاده می‌شود (به عنوان مثال Kisi & Shiri, 2012). به منظور شبیه‌سازی، از ۱۶۵ داده، ۱۱۵ داده (حدود ۷۰ درصد) برای آموزش و ۵۰ داده (حدود ۳۰ درصد) برای آزمایش به‌صورت تصادفی در نظر گرفته شدند. علت این کار این است که هر دو مجموعه آموزش و آزمایش، ترکیبی نماینده از تمامی شرایط هیدرولوژیکی (کم‌آبی، نرمال و پربابی) داشته باشند و ارزیابی مدل‌ها از نظر آماری کمتر دچار بایاس ناشی از ترتیب زمانی شود (Darabi

¹ Interquartile Range

Seyedian et al. 2017; AIDahoul et al. 2022; Moradinejad 2024) و انجام آزمون حساسیت بر روی داده‌های ایستگاه آستانه رودخانه سفیدرود تعیین شده‌اند، به‌طوری که برای هر مدل، پیکربندی نهایی، کمترین RMSE و بیشترین R^2 را در داده‌های آزمایش بدون نشانه‌های بیش‌برازش فراهم کرده است. پس از شبیه‌سازی، نتایج با شاخص‌های آماری ضریب تعیین (R^2)، میانگین مربعات خطا (MSE) و ریشه میانگین مربعات خطا (RMSE) اعتبارسنجی شده است. این شاخص‌ها امکان مقایسه سطح دقت و قدرت پیش‌بینی مدل‌های مختلف را فراهم می‌کنند و در انتخاب مدل برتر نقش مهمی دارند (Kisi & Shiri, 2012; Zangeneh, 2024; Asadi et al., 2025; Bezak et al., 2024).

۳- نتایج و بحث

پس از مدل‌سازی، نتایج آماری حاصله در جدول ۱ (مقادیر بدون نرمال‌سازی) و همچنین نتایج نموداری حاصل از شبیه‌سازی رودخانه سفیدرود در قالب شکل‌های ۲ تا ۶ ارائه شده است. همانطور که در شکل ۲ مشاهده می‌شود، قسمت بالایی این شکل مربوط به نمودارهای سری زمانی رسوب معلق شبیه‌سازی شده و مشاهداتی به ازای شش مدل برای داده‌های آموزش و قسمت پایینی آن مربوط به نمودارهای داده آموزش رسوب معلق شبیه‌سازی شده در مقابل رسوب مشاهداتی برای رودخانه سفیدرود است. همچنین، شکل ۳ نیز مربوط به نتایج داده‌های آزمایش شش مدل مذکور برای رودخانه سفیدرود است.

جدول ۱. نتایج آماری داده‌های آموزش و آزمایش رودخانه سفیدرود براساس شاخص‌های آماری

Table 1. Statistical results of the Sefidrood River training and testing data based on statistical indices.

Training			Testing			مدل	رودخانه
RMSE	MSE	R^2	RMSE	MSE	R^2		
۴/۲	۱۷/۲۶	۰/۸۵	۴/۴۵	۱۹/۷۷	۰/۸۴	ANN	سفیدرود
۴/۲	۱۷/۶۱	۰/۸۷	۴/۶۱	۲۱/۲۸	۰/۸۶	ANN-GA	
۴/۴۲	۱۹/۵	۰/۸۴	۴/۱۴	۱۷/۱۶	۰/۸۶	ANN-ICA	
۵/۳۸	۲۸/۹۹	۰/۷۶	۵/۵	۳۰/۲	۰/۷۶	ANN-PSO	
۳/۴۹	۱۲/۱۵	۰/۹۰	۴/۶۰	۲۱/۱۷	۰/۸۳	Random Forest	
۴/۱۳	۱۷/۱	۰/۸۶	۴/۳۷	۱۹/۰۶	۰/۸۵	SVM	

با توجه به شکل‌های ۲ و ۳ و جدول ۱، مدل RF بهترین عملکرد را در مرحله آموزش با $R^2 = 0.90$ و $RMSE = 3.49$ داشته است. همچنین این مدل ضمن تقریب بسیار خوب آن در داده‌های آموزش، باعث تجمع نمودارهای جفتی شبیه‌سازی-مشاهده حول خط ۱:۱ شده است. علاوه بر آن، مدل‌های ANN و SVM نیز عملکرد پایدار دارند ($R^2_{ANN} = 0.85$ و $R^2_{SVM} = 0.86$). هر دو مدل تخمینی‌های نزدیک به واقعیت دارند و کمتر دچار خطای سیستماتیک هستند. در میان مدل‌های هیبریدی، ANN-GA در مرحله آموزش بهترین

لایه پنهان ثابت نگه داشته شد. در نسخه ANN-PSO، مدل ازدحام ذرات برای تنظیم بردارهای وزن و بایاس به کار گرفته شد و ذرات در فضای جست‌وجو حول بهترین موقعیت فردی و جمعی حرکت کردند (Baazi et al., 2014; Azar et al., 2021). تعداد ذرات برابر ۵۰، تعداد تکرار ۲۰۰ و ضرایب شناختی و اجتماعی هر دو برابر ۱/۴۹۶ انتخاب شد تا تعادلی بین جست‌وجوی محلی و سراسری برقرار شود، در حالی که معماری پایه ANN همان ساختار یک‌لایه با ۱۰ نورون بود. در مدل ANN-ICA، مدل رقابت استعماری برای یافتن ترکیب بهینه وزن‌ها و پارامترهای شبکه به کار رفت و کشورها (جواب‌های کاندید) در قالب امپراتوری‌های مختلف سازمان‌دهی شدند (Emami et al., 2018; Choopan, 2019). در این طرح، تعداد کل کشورها ۱۰۰، تعداد امپریالیست‌ها ۳۰، و تعداد تکرارها ۱۵۰ در نظر گرفته شد و ضرایب $\beta = 2$ و $\zeta = 0.02$ برای کنترل شدت جذب و انقلاب تنظیم شد، در حالی که تعداد نورون‌های لایه پنهان ANN برابر ۱۰ ثابت ماند.

۲-۳-۳- مدل‌های یادگیری ماشین

جنگل تصادفی (RF) به عنوان یک روش دسته‌جمعی مبتنی بر مجموعه‌ای از درخت‌های تصمیم، با نمونه‌برداری بوت‌استرپ^۱ (با همان نمونه‌گیری با جایگذاری) از داده‌ها و انتخاب تصادفی زیرمجموعه‌ای از ویژگی‌ها در هر گره برای مدل‌سازی ارتباط غیرخطی بین دبی و رسوب معلق به کار گرفته شد (Seyedian et al., 2017; AIDahoul et al., 2022). در این مطالعه، تعداد درخت‌ها ۲۰۰ تعیین شد و بیشینه عمق درخت‌ها بدون محدودیت، حداقل تعداد نمونه برای تقسیم گره ۱۰ و حداقل تعداد نمونه در برگ ۱ در نظر گرفته شد تا تعادل مناسبی بین دقت و پیچیدگی مدل ایجاد شود. ماشین بردار پشتیبان (SVM) نیز با استفاده از کرنل شعاعی (RBF) برای تقریب نگاشت غیرخطی داده‌های ورودی به فضای ویژگی با بعد بالاتر مورد استفاده قرار گرفت (Latif et al., 2023; Le et al., 2024). پارامتر جریمه C برابر ۱ انتخاب شد، تابع هسته‌ای از نوع RBF، پارامتر Gamma براساس تنظیم Auto و آستانه حساسیت خطای مجاز برابر ۰/۱ در نظر گرفته شد تا ضمن کنترل میزان انحراف مجاز، پایداری مدل در برابر داده‌های پرت حفظ شود.

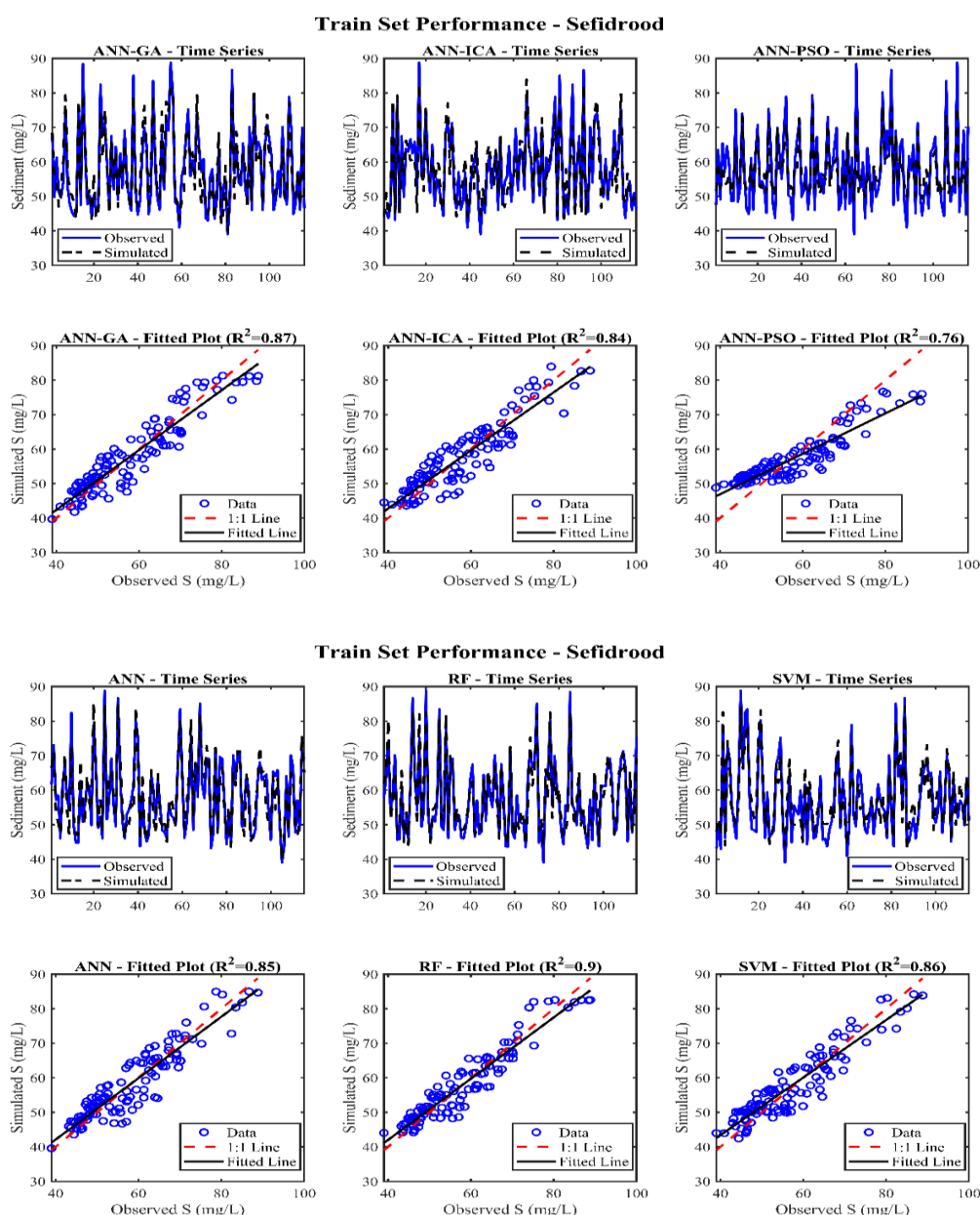
۲-۴- واسنجی مدل

همانطور که در بخش‌های قبل ذکر شد، مدل‌سازی براساس تمامی پارامترهای عددی انتخاب‌شده برای ساختار ANN، مدل‌های فراابتکاری و مدل‌های RF و SVM انجام شده و این انتخاب پارامترها براساس ترکیبی از توصیه‌های مطالعات پیشین در مدل‌سازی رسوب و کیفیت آب (Razavizadeh & Dargahian 2018; Kisi & Shiri 2012; Sanikhani et al., 2012; Naseri et al. 2019; Nikoo et al. 2018; Emami et al. 2018; Baazi et al. 2014;

¹ Bootstrap Sampling

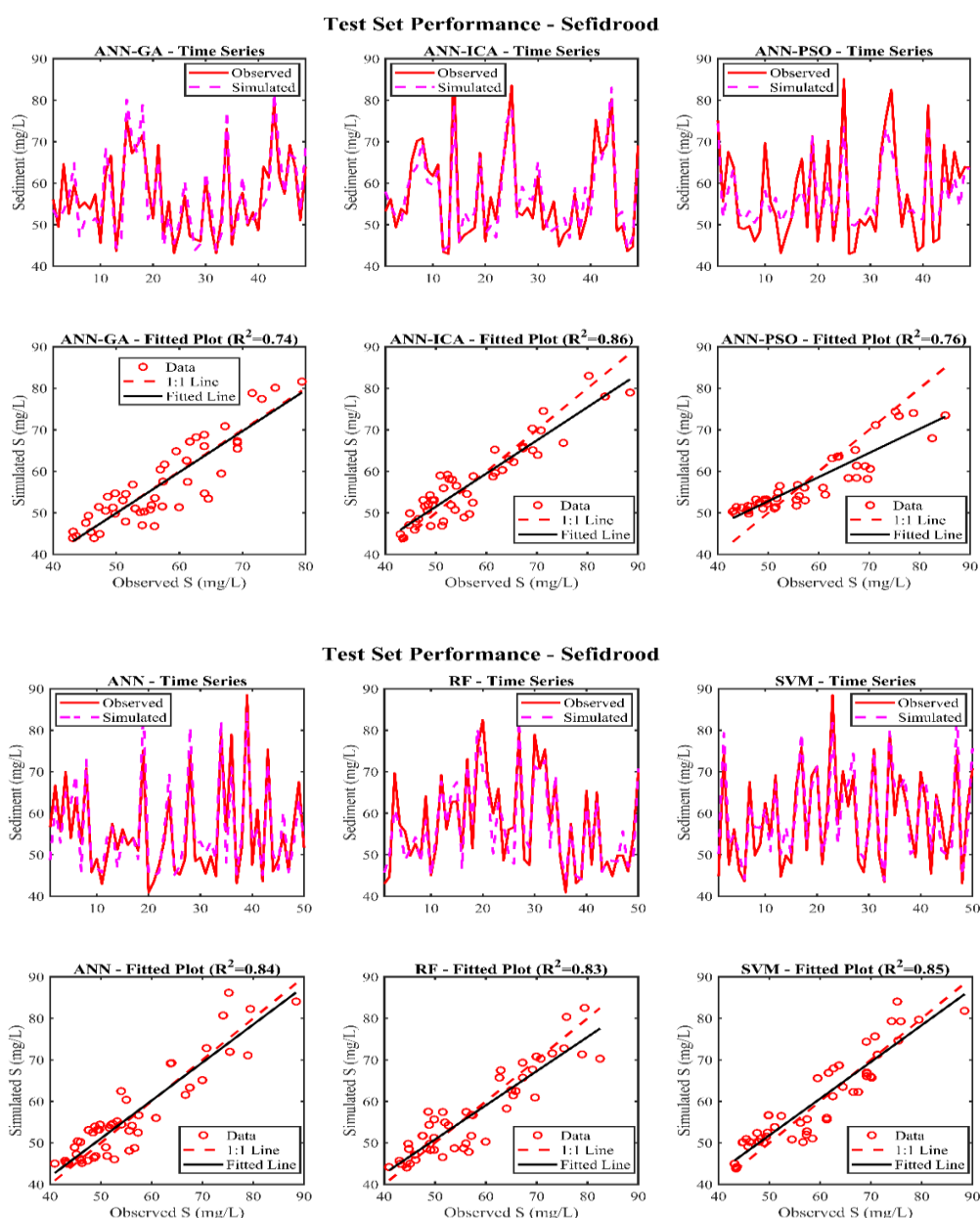
عملکرد را ارائه کرد ($R^2 = 0.87$)، در حالی که سایر مدل‌های ترکیبی ANN-PSO و ANN-ICA تفاوت عملکردی قابل توجهی نسبت به ANN پایه نشان ندادند. با این حال، ارزیابی مدل‌ها در مرحله آزمایش مستقل تفاوت‌هایی در قابلیت تعمیم آن‌ها آشکار کرد. در این مرحله، مدل ANN-ICA با کمترین مقدار $RMSE = 4.14$ و مقدار بالای $R^2 = 0.86$ بهترین عملکرد کلی را نشان داد. همچنین، مدل ANN-GA با $R^2 = 0.86$ عملکرد مشابهی داشت، در حالی که مدل SVM با $R^2 = 0.85$ و $RMSE = 4.37$ نیز دقت قابل قبولی ارائه کرد. در مقابل، مدل RF که در مرحله آموزش عملکرد بسیار مطلوبی داشت، در مرحله آزمایش با کاهش نسبی عملکرد ($R^2 =$

0.83) مواجه شد که می‌تواند نشان‌دهنده کاهش قابلیت تعمیم آن نسبت به برخی مدل‌های هیبریدی باشد. این نتایج با مطالعات پیشین در زمینه پیش‌بینی بار معلق رسوبی که عملکرد رقابتی روش‌های یادگیری ماشین، به‌ویژه SVM و مدل‌های هیبریدی، را در مواجهه با داده‌های غیرخطی و پرنوسان گزارش کرده‌اند، همخوانی دارد. برای مثال، Essam et al. (2022) و Allawi et al. (2023) توانایی بالای مدل‌های داده‌محور در شبیه‌سازی بار معلق را نشان داده‌اند. در مطالعه حاضر نیز مدل‌های هیبریدی ANN، به‌ویژه ANN-ICA، توانایی مناسبی در استخراج روابط غیرخطی بین دبی جریان و بار رسوب معلق رودخانه سفیدرود نشان دادند.



شکل ۲. نمودارهای سری زمانی داده‌های آموزش رسوب معلق شبیه‌سازی شده و مشاهداتی به ازای شش مدل و همین‌طور نمودارهای داده آموزش رسوب معلق شبیه‌سازی شده در مقابل رسوب مشاهداتی برای رودخانه سفیدرود

Fig 2. Time series plots of simulated and observed suspended sediment training data for the six models, as well as plots of simulated suspended sediment training data versus observed sediment for the Sefidrood River.



شکل ۳. نمودارهای سری زمانی داده‌های آزمایش رسوب معلق شبیه‌سازی شده و مشاهداتی به ازای شش مدل و همین‌طور نمودارهای داده آزمایش رسوب معلق شبیه‌سازی شده در مقابل رسوب مشاهداتی برای رودخانه سفیدرود

Fig 3. Time series plots of simulated and observed suspended sediment test data for the six models, as well as plots of simulated suspended sediment test data versus observed sediment for the Sefidrood River.

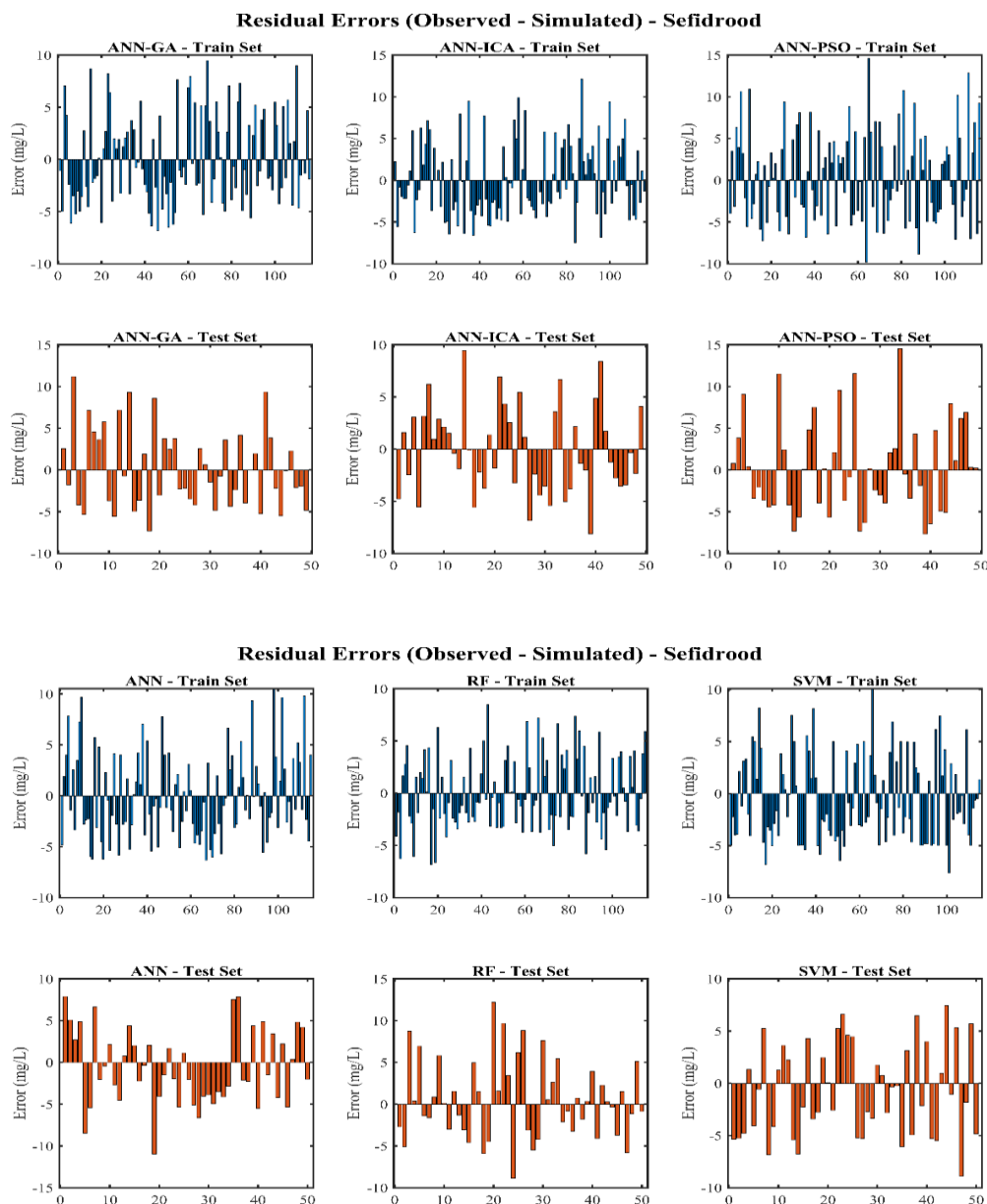
در کار (2011) Melesse et al. نیز ذکر شده است.

با توجه به شکل ۴، توزیع خطاهای رودخانه سفیدرود در کل مدل‌ها حول محور صفر و با پراکندگی معقول است. این توزیع در مدل‌های ANN، ANN-ICA، و RF متراکم‌تر و نزدیک‌تر به محور صفر مشاهده شد. در داده‌های آزمایش، این توزیع خطا بین ۱۰- تا ۱۰ میلی گرم (به جز مدل ANN-PSO) دیده می‌شود و این امر مطابق با مطالعات دیگر مانند Kisi (2007) در داده‌های هیدرولوژیکی گاهی نوسانات شدید پذیرفته می‌شود. در مدل ANN-ICA در داده‌های آزمایش (با R^2 بالاتر و RMSE کمتر) توزیع

به‌طور کلی، مقدار R^2 و RMSE بدست آمده در این پژوهش، اغلب مقالات معتبر حوضه کیفیت آب و رسوب سازگاری دارد، بطور مثال در پژوهش Kisi & Ozkan (2017) مقادیر R^2 برای RF و ANN در تخمین رسوب معمولاً در بازه ۰/۸ تا ۰/۹ گزارش شده است. همچنین، تفاوت زیاد مدل ترکیبی ANN با GA و PSO نسبت به ICA، بیانگر این است که همه‌ی مدل‌های فراابتکاری الزاماً مدل را بهتر نمی‌کنند. ANN-ICA تنها مدلی بود که در مرحله آزمایش مقدار R^2 بالاتری نسبت به مرحله آموزش نشان داد، که نشان‌دهنده پایداری مناسب عملکرد آن در داده‌های مستقل است، که این ویژگی

مرحله آزمایش نشان داد که تعمیم‌پذیری آن کمتر از برخی مدل‌های هیبریدی بوده است، در حالی که مدل SVM در برخی بازه‌ها پراکندگی بیشتری در خطاها نشان داد، اگرچه میانگین خطای آن همچنان نزدیک صفر باقی ماند.

خطای کمتری دارد و ارزیابی کیفی این مدل را تقویت می‌کند. بالعکس، مدل‌های فراابتکاری GA و PSO پراکندگی بیشتری در خطاهای آزمایش دارند. مدل RF در مرحله آموزش خطاهای پراکنده (یا مجزا) کمتری نشان داد، اگرچه کاهش نسبی عملکرد آن در



شکل ۴. نمودارهای خطاهای باقی‌مانده داده‌های آزمایش و آموزش برای رودخانه سفیدرود
Fig 4. Residual error plots of the test and training data for the Sefidrood River

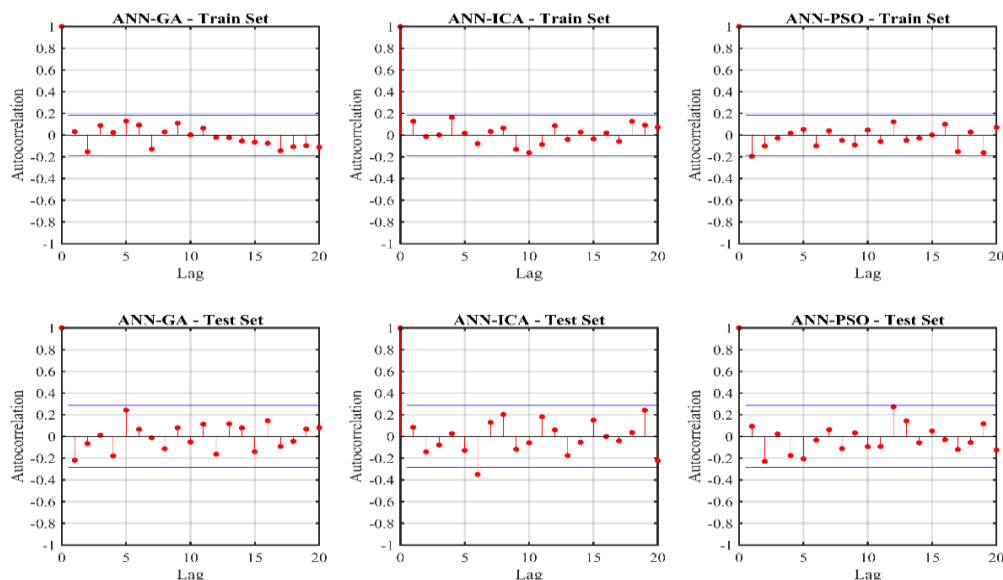
(2017)، (2023) Allawi et al. و (2020) Sharafati et al. که نشان دادند، عدم وجود همبستگی شدید در خطاها برای مدل‌های زمانی خوب ضروری است و بیانگر عدم وجود ساختار سیستماتیک در باقی‌مانده‌ها است. در داده‌های آزمایش، کمی نوسان بیشتر در بازه‌های پایین دیده می‌شود، اما به هیچ وجه همبستگی معناداری وجود ندارد، یعنی این رفتار نشان می‌دهد که مدل‌ها توانسته‌اند الگوهای اصلی تغییرات زمانی بار رسوب را بدون ایجاد ساختار

نمودارهای شکل ۵ بیانگر همبستگی خطاهای داده‌های آزمایش و آموزش برای رودخانه سفیدرود است. برای رودخانه سفیدرود، همه مدل‌ها (در هر دو مجموعه داده) دارای باقی‌مانده‌هایی با همبستگی بسیار نزدیک به صفر در اکثر بازه‌ها هستند. این نقاط عمدتاً بین خطوط آبی (حدود اطمینان ۹۵ درصد) قرار گرفته‌اند و در برخی بازه‌ها، لندگی انحراف دارند که برحسب اتفاق و معمولاً ناشی از محدودیت حجم داده‌ها است، مانند کارهای Kisi & Ozkan

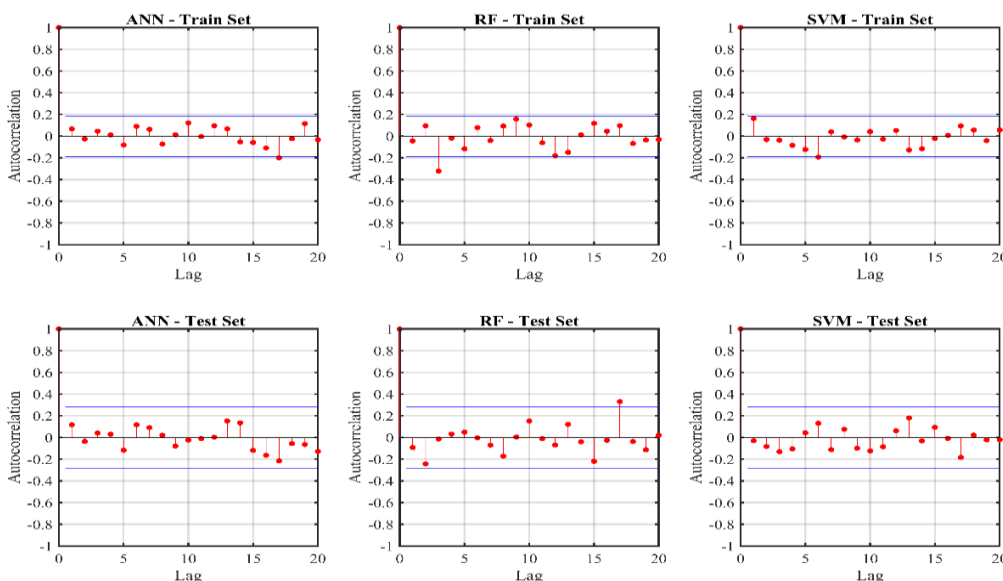
دارد از نظر آماری قلیل اعتماد گزارش شده‌اند. بنابراین، نمودارهای حاضر این معیار مهم را برآورده می‌کنند.

مشخص در باقی‌مانده‌ها بازنمایی کنند. در پژوهش Noori et al. (2012) تنها مدل‌هایی که خطای آن‌ها رفتار تصادفی و غیر وابسته

Autocorrelation of Residual Errors - Sefidrood



Autocorrelation of Residual Errors - Sefidrood



شکل ۵. خودهمبستگی نگاره‌های خطاهای مدل در بخش‌های مربوط به داده‌های آزمایش و آموزش برای رودخانه سفیدرود

Fig 5. Autocorrelation of model error graphs in the sections related to testing and training data for the Sefidrood River

نشان‌دهنده احتمال بیش‌برازش^۱ و حساسیت به داده‌های پرت است. مقایسه نتایج این مطالعه با تحقیقات قبلی نیز تأیید می‌کند که مدل‌های هیبریدی و RF در برآورد رسوب معلق رودخانه‌ای می‌توانند ابزارهای مؤثری برای کاهش خطای شبیه‌سازی رسوب معلق باشند. (Allawi et al., 2023). نرمال بودن خطای باقیمانده برای اعتبار مدل‌سازی رسوب معلق ضروری است، زیرا عدم نرمالیتی می‌تواند

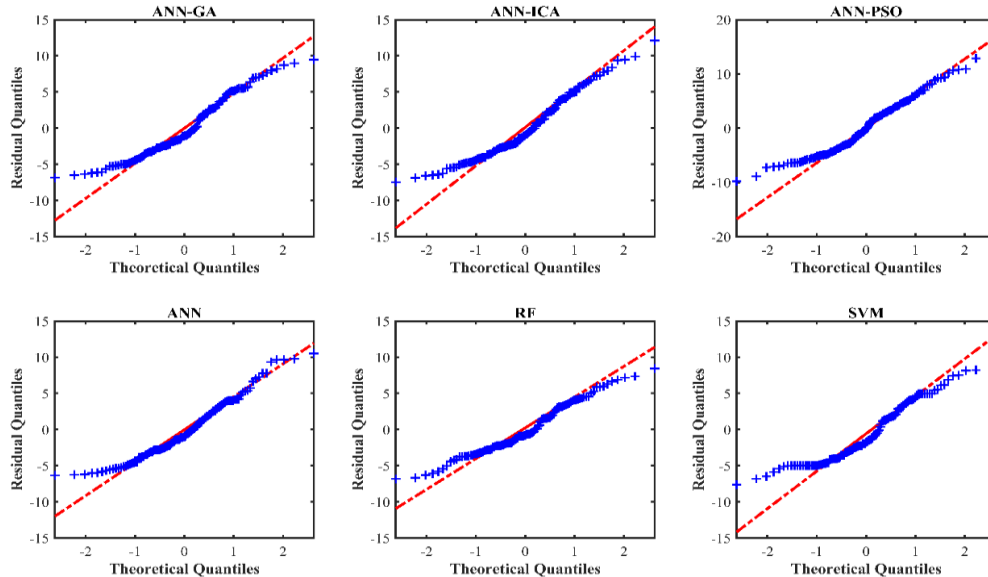
شکل ۶، نمودارهای QQ برای مدل‌های مختلف نشان می‌دهد که مدل‌های ANN-GA، ANN-ICA و RF در هر دو مرحله آموزش و آزمایش انحراف کمتری از توزیع نرمال در باقی‌مانده‌ها نشان دادند که بیانگر رفتار مناسب خطاهای آن‌ها است و از نظر پایداری و دقت عملکرد بهتری نسبت به سایر مدل‌ها ارائه می‌دهند. در مقابل، مدل ANN-PSO بیشترین انحراف از خط نرمال را نشان می‌دهد که

¹ Overfitting

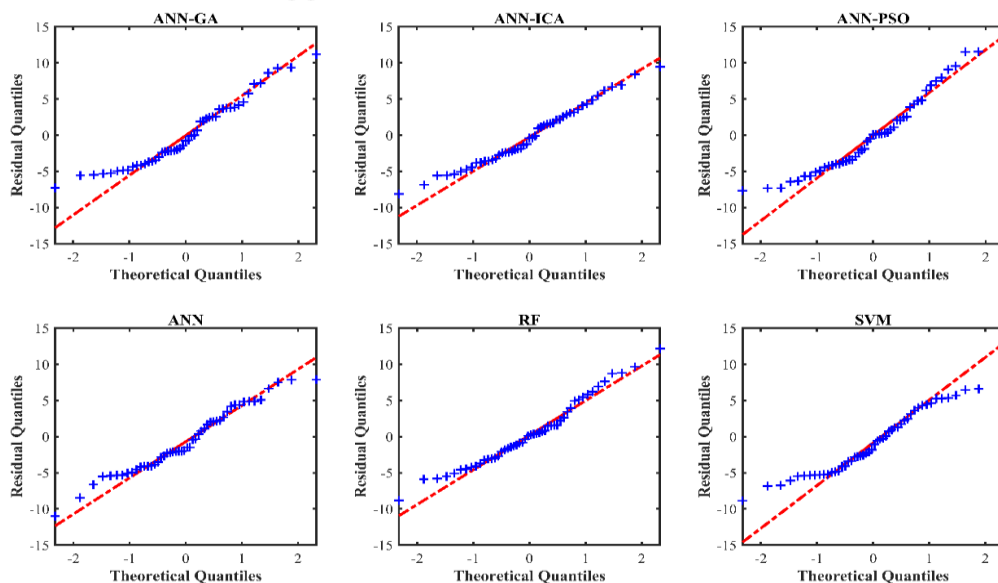
در کاهش انحرافات و بهبود نرمالیتی خطای باقیمانده‌ها تأیید کرده‌اند.

منجر به تخمین‌های سوگیر شود. نتایج این مطالعه با مطالعات دیگر مانند (Allawi et al. (2023)، Li & Aldahoul et al. (2021) و Yang (2022) همخوانی دارد که برتری مدل‌های هیبریدی و RF را

QQ Plot of Train Residual Errors - Sefidrood



QQ Plot of Test Residual Errors - Sefidrood



شکل ۶. نمودارهای QQ در مقابل توزیع خطای باقیمانده در داده‌های آزمایش و آموزش برای رودخانه سفیدرود

Fig 6. QQ plots versus residual error distribution in test and training data for the Sefidrood River

این آزمون به‌ویژه زمانی توصیه می‌شود که چند مدل مختلف بر روی یک مجموعه داده یکسان ارزیابی شده باشند و فرض نرمال بودن خطاها برقرار نباشد. در این پژوهش، قدرمطلق خطای پیش‌بینی داده‌های آزمایش برای شش مدل ANN-PSO، ANN-GA، ANN-ICA، RF و SVM محاسبه و ماتریس حاصل به عنوان ورودی آزمون فریدمن مورد استفاده قرار گرفت.

علاوه بر شاخص‌های متداول ارزیابی عملکرد شامل ضریب تعیین (R^2)، میانگین مربعات خطا (MSE) و ریشه میانگین مربعات خطا (RMSE) (جدول ۱)، به‌منظور بررسی اینکه آیا اختلاف عملکرد مدل‌های مورد مطالعه از نظر آماری معنادار است یا خیر، از آزمون ناپارامتری فریدمن استفاده شد. این آزمون برای مقایسه چند تیمار وابسته و رتبه‌بندی آن‌ها براساس میزان خطا به کار می‌رود و معادل ناپارامتری تحلیل واریانس با اندازه‌گیری‌های تکراری است. استفاده از

بین مدل ANN و مدل‌های ANN-GA ($p = 0.000023$)، ANN-ICA ($p = 0.000018$)، JCA ($p = 0.0000016$)، RF ($p = 0.00000016$) و در نهایت SVM ($p = 0.000018$) از نظر آماری معنادار است. در مقابل، اختلاف بین ANN و ANN-PSO در مرز سطح معناداری قرار داشت ($p=0.053$) و از نظر آماری معنادار محسوب نشد.

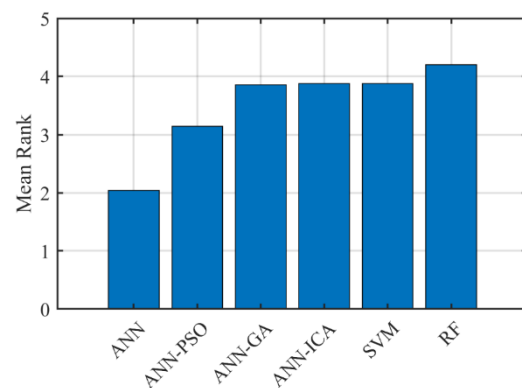
جدول ۲. نتایج آزمون تعقیبی بونفرونی برای مقایسه زوجی عملکرد مدل‌ها

Table 2. Bonferroni post-hoc test results for pairwise comparison of model performance

pValue	UpperCI	MeanDifference	LowerCI	مدل دوم	مدل اول
۰/۰۰۰۰۲۳	-۰/۷۰۷	-۱/۸۱۶۳	-۲/۹۲۵۷	ANN-GA	ANN
۰/۰۵۳	-۰/۰۰۷۴	-۱/۱۰۲	-۲/۲۱۱۴	ANN-PSO	ANN
۰/۰۰۰۰۱۸	-۰/۷۳۷۳	-۱/۸۳۷	-۲/۹۴۶۱	ANN-ICA	ANN
۰/۰۰۰۰۰۱۶	-۱/۰۵۴	-۲/۱۶۳	-۳/۲۷۳۷	RF	ANN
۰/۰۰۰۰۱۸	-۰/۷۳	-۱/۸۴	-۲/۹۴۶۱	SVM	ANN
۰/۸۸۱۷	۱/۸۳۴	۰/۷۱۴۳	-۰/۳۹۵۱	ANN-PSO	ANN-GA
۱	۱/۰۸۹	-۰/۰۲۰۴	-۱/۱۲۹۸	ANN-ICA	ANN-GA
۱	۰/۷۶۳	-۰/۳۴۶۹	-۱/۴۵۶۳	RF	ANN-GA
۱	۱/۰۸۹	-۰/۰۲۰۴	-۱/۱۲۹۸	SVM	ANN-GA
۰/۷۷۹	۰/۳۷۵	-۰/۷۳۴۷	-۱/۸۴۴۱	ANN-ICA	ANN-PSO
۰/۰۷۴	۰/۰۴۸۲	-۱/۰۶۱۲	-۲/۱۷۰۶	RF	ANN-PSO
۰/۷۸	۱/۳۷۵	-۰/۷۳۴۷	-۱/۸۴۴۱	SVM	ANN-PSO
۱	۰/۷۸۳	-۰/۳۲۶۵	-۱/۴۳۵۹	RF	ANN-ICA
۱	۱/۱۰۹۴	۰	-۱/۱۰۹۴	SVM	ANN-ICA
۱	۱/۴۳۶	۰/۳۲۶۵	-۰/۷۸۲۹	SVM	RF

همچنین، برای تمامی مقایسه‌های بین مدل‌های ANN-GA، ANN-PSO، ANN-ICA، SVM و RF، مقادیر احتمال بزرگ‌تر از ۰/۰۵ به دست آمد؛ بنابراین، اختلاف عملکرد این مدل‌ها از نظر آماری معنادار نبود. این نتیجه نشان می‌دهد که اگرچه آزمون فریدمن مدل ANN را از نظر رتبه خطا در جایگاهی متفاوت نسبت به برخی مدل‌ها قرار می‌دهد، اما بین مدل‌های ANN-PSO، ANN-GA، ANN-ICA، SVM و RF اختلاف آماری قابل اثباتی وجود ندارد و این مدل‌ها را می‌توان از دیدگاه آماری در یک گروه عملکردی مشابه قرار داد. از این رو، انتخاب نهایی میان این مدل‌ها بهتر است علاوه بر نتایج آزمون آماری، براساس شاخص‌های دقت (R^2 و RMSE)، توان تعمیم‌پذیری مدل، هزینه محاسباتی و الزامات کاربردی مسئله انجام شود. در مجموع، آزمون فریدمن وجود اختلاف معنادار بین عملکرد مدل‌های مورد بررسی را تأیید کرد، در حالی که آزمون تعقیبی بونفرونی نشان داد این اختلاف عمدتاً ناشی از تفاوت مدل ANN با برخی از مدل‌های دیگر است و بین مدل‌های پیشرفته‌تر شامل ANN-GA، ANN-PSO، ANN-ICA، SVM و RF اختلاف معناداری مشاهده نمی‌شود. بنابراین، نتایج آماری این پژوهش در کنار شاخص‌های RMSE و R^2 نشان می‌دهد که ارزیابی عملکرد مدل‌ها باید بر پایه مجموعه‌ای از معیارهای آماری انجام شود و اتکا به یک شاخص منفرد نمی‌تواند تصویری کامل از کیفیت پیش‌بینی

نتایج آزمون فریدمن برای داده‌های رودخانه سفیدرود نشان داد که اختلاف عملکرد مدل‌ها از نظر آماری کاملاً معنادار است ($p = 2.0 \times 10^{-8} < 0.05$). بنابراین، فرض صفر مبنی بر یکسان بودن عملکرد همه مدل‌ها رد شد و مشخص شد که حداقل یک مدل از نظر توزیع خطاهای پیش‌بینی، اختلاف معناداری با سایر مدل‌ها دارد. برای تفسیر دقیق‌تر نتایج، میانگین رتبه خطا برای هر مدل محاسبه شد (شکل ۷). از آنجا که در آزمون فریدمن میانگین رتبه کمتر نشان‌دهنده عملکرد بهتر مدل است، مدل ANN با میانگین رتبه ۲/۰۴ بهترین عملکرد را به خود اختصاص داد. پس از آن، مدل‌های ANN-PSO (۳/۱۴)، ANN-GA (۳/۸۶)، ANN-ICA (۳/۸۸)، SVM (۳/۸۸) و در نهایت RF با میانگین رتبه ۴/۲۰ قرار گرفتند.



شکل ۷- میانگین رتبه مدل‌های مورد بررسی براساس آزمون فریدمن

Fig 7- Mean ranks obtained for the evaluated models using the Friedman test

در نگاه نخست، این رتبه‌بندی ممکن است با نتایج شاخص‌های کلی ارائه‌شده در جدول ۱ متفاوت به نظر برسد؛ زیرا براساس مقادیر R^2 و RMSE، مدل RF در مرحله آموزش و مدل‌های ANN-ICA و SVM در مرحله آزمایش عملکرد مطلوب‌تری نشان داده‌اند. با این حال، این تفاوت ناشی از ماهیت متفاوت معیارهای ارزیابی است. شاخص‌هایی مانند R^2 و RMSE تنها یک مقدار خلاصه از کل عملکرد مدل ارائه می‌کنند، در حالی که آزمون فریدمن بر مبنای رتبه‌بندی خطای هر نمونه در تمام گام‌های زمانی انجام می‌شود و رفتار کامل سری زمانی خطا را مورد ارزیابی قرار می‌دهد. در نتیجه، ممکن است مدلی دارای RMSE کلی پایین باشد، اما در تعدادی از مشاهدات خطاهای بزرگ‌تری ایجاد کند و در رتبه‌بندی فریدمن جایگاه پایین‌تری کسب نماید. بنابراین، اختلاف مشاهده‌شده میان نتایج شاخص‌های کلاسیک و آزمون فریدمن لزوماً بیانگر تناقض نیست، بلکه ناشی از تفاوت مبنای آماری این روش‌های ارزیابی است.

به‌منظور تعیین زوج مدل‌هایی که اختلاف آن‌ها از نظر آماری معنادار است، پس از آزمون فریدمن از آزمون مقایسه‌های چندگانه با اصلاح بونفرونی استفاده شد (جدول ۲). نتایج نشان داد که اختلاف

مدل‌ها ارائه دهد.

داده‌های این پژوهش انجام شده و کاربرد آن‌ها در سایر حوضه‌ها ممکن است نیازمند بازکالیبراسیون باشد.

۴- نتیجه‌گیری

در این پژوهش، شش مدل داده‌محور شامل ANN، سه نسخه هیبرید فراابتکاری آن (ANN-ICA، ANN-PSO، ANN-GA) و دو الگوریتم یادگیری ماشین (RF و SVM) برای برآورد بار معلق رسوب رودخانه سفیدرود ارزیابی شدند. نتایج نشان داد که مدل‌های پیشرفته یادگیری ماشین، به‌ویژه RF و SVM، براساس شاخص‌های R^2 و RMSE دقت بالایی در مراحل آموزش و آزمایش داشتند، در حالی که مدل هیبرید ANN-ICA از نظر تعمیم‌پذیری و مقاومت در برابر بیش‌برازش عملکرد مطلوبی در داده‌های آزمایش نشان داد. به‌طور کلی، هر یک از مدل‌های بررسی‌شده دارای نقاط قوت متفاوتی بودند؛ به گونه‌ای که RF و SVM از نظر شاخص‌های کلی دقت و ANN-ICA از نظر قابلیت تعمیم عملکرد مناسبی ارائه کردند. ازاین‌رو، استفاده از این مدل‌ها می‌تواند ابزار قابل‌اعتمادی برای برآورد روزانه رسوبات معلق و پشتیبانی از مدیریت رسوب در حوضه سفیدرود فراهم کند. علاوه بر شاخص‌های عملکرد، آزمون ناپارامتری فریدمن نیز وجود اختلاف آماری معنادار بین عملکرد مدل‌ها را تأیید کرد ($p < 0.05$). همچنین، نتایج آزمون تعقیبی بونفرونی نشان داد که اختلاف معنادار عمدتاً مربوط به مقایسه مدل ANN با برخی از مدل‌های دیگر بوده، در حالی که بین مدل‌های ANN-PSO، ANN-GA، ANN-ICA، SVM و RF اختلاف آماری معناداری مشاهده نشد. این نتایج نشان می‌دهد که ارزیابی عملکرد مدل‌ها باید بر پایه ترکیبی از شاخص‌های آماری (R^2 و RMSE) و آزمون‌های استنباطی انجام شود و اتکا به یک معیار منفرد نمی‌تواند تصویر کاملی از رفتار مدل‌ها ارائه دهد.

به‌کارگیری این مدل‌ها در سامانه‌های بهره‌برداری و پایش می‌تواند به برنامه‌ریزی برای کاهش رسوبات ورودی به مخازن، کنترل آلودگی رسوب‌محور و پشتیبانی از تصمیم‌گیری در طرح‌های انتقال و تأمین آب حوضه سفیدرود کمک کند. همچنین، نتایج این پژوهش نشان داد که در محیط‌های پیچیده اقلیمی-هیدرولوژیکی استان گیلان، مدل‌های داده‌محور پیشرفته نسبت به روش‌های کلاسیک رگرسیونی ظرفیت بیشتری برای شبیه‌سازی روابط غیرخطی و برآورد رسوب معلق دارند. این مطالعه بر پایه داده‌های روزانه دبی و بار معلق رسوب یک ایستگاه منتخب (آستانه) در رودخانه سفیدرود طی دوره زمانی ۱۳۸۶ تا ۱۴۰۱ انجام شد؛ بنابراین، تعمیم نتایج به سایر رودخانه‌ها یا کل حوضه باید با احتیاط و در چارچوب شرایط هیدرولوژیکی مشابه صورت گیرد. همچنین، تنها متغیرهای اصلی هیدرولوژیکی (دبی و بار معلق رسوب) به عنوان ورودی مدل‌ها استفاده شدند و اثر عواملی مانند بارندگی، کاربری اراضی، دما و تغییرات مورفولوژیکی کانال به‌طور مستقیم در مدل لحاظ نشد که می‌تواند بخشی از عدم قطعیت نتایج را توضیح دهد. افزون بر این، تنظیم پارامترهای مدل‌ها براساس

منابع

- Aires, U. R. V., da Silva, D. D., Fernandes Filho, E. I., Rodrigues, L. N., Uliana, E. M., Amorim, R. S. S., ... and Campos, J. A. 2023. Machine learning-based modeling of surface sediment concentration in Doce river basin. *Journal of Hydrology*, 619, 129320. doi:10.1016/j.jhydrol.2023.129320
- AlDahoul, N., Ahmed, A.N., Allawi, M.F., Sherif, M., Sefelnasr, A., Chau, K.W., and El-Shafie, A. 2022. A comparison of machine learning models for suspended sediment load classification. *Engineering Applications of Computational Fluid Mechanics*, 16(1), 1211-1232. doi:10.1080/19942060.2022.2073565
- AlDahoul, N., Essam, Y., Kumar, P., Ahmed, A.N., Abdullatif, M., El-Shafie, A., Sherif, M., and Sefelnasr, A. 2021. Suspended sediment load prediction using long short-term memory neural network. *Scientific Reports*, 11, 7826. doi:10.1038/s41598-021-87415-4
- Allawi, M.F., Sulaiman, S.O., Sayl, K.N., Sherif, M., and El-Shafie, A. 2023. Suspended sediment load prediction modelling based on artificial intelligence methods: The tropical region as a case study. *Heliyon*, 9(8), e18506. doi:10.j.heliyon.2023.e18506
- Azari, A., Soori, S., and Bonakdari, H. 2017. The Application of Imperialist Competitive Algorithm in Determining the Optimal Parameters of Empirical Area Reduction Method to Predict the Sedimentation Process in Dez Dam. *Geography and Environmental Sustainability*, 7(3), 1-14.
- Baazi, H., Kalla, M., and Tebbi, F.Z. 2014. PSO-ANN's based suspended sediment concentration in Ksob basin, Algeria. *Journal of Engineering and Technology Research*, 6(5), 98-107. doi:10.5897/JETR2014.0549

- Essam, Y., Huang, Y.F., Birima, A.H. et al. 2022. Predicting suspended sediment load in Peninsular Malaysia using support vector machine and deep learning algorithms. *Sci Rep* 12, 302. doi:10.1038/s41598-021-04419-w
- Fardadi Shilsar, M.J., Mazaheri, M., and Mohammad Vali Samani, J. 2021. Analytical solution of the pollutant transport equation with variable coefficients in rivers using the Laplace transform. *Water and Irrigation Management*, 11(4), 683–698. doi:10.22059/jwim.2021.329149.911 [In Persian]
- Fardadi Shilsar, M.J., Mazaheri, M., and Mohammad Vali Samani, J. 2022a. Analytical solution of mass transport equation in river networks. *Sharif Journal of Mechanical Engineering*, 38.3(1), 35–49. doi:10.24200/j40.2021.58767.1613 [In Persian]
- Fardadi Shilsar, M.J., Mazaheri, M., and Mohammad Vali Samani, J. 2022b. Analytical solution of pollutant transport equation in different types of river networks considering distributed source term. *Iranian Journal of Soil and Water Research*, 53(5), 1057–1077. doi:10.22059/ijswr.2022.341884.669250 [In Persian]
- Fardadi Shilsar, M.J., Mazaheri, M., and Mohammad Vali Samani, J. 2023. A semi-analytical solution for one-dimensional pollutant transport equation in different types of river networks. *Journal of Hydrology*, 619, 129287. doi:10.1016/j.jhydrol.2023.129287
- Fawell, J., and Nieuwenhuijsen, M.J. 2003. Contaminants in drinking water: Environmental pollution and health. *British Medical Bulletin*, 68(1), 199-208. doi:10.1093/bmb/ldg027
- Gul, E., Safari, M.J.S., Torabi Haghighi, A., and Danandeh Mehr, A. 2021. Sediment transport modeling in open channels using machine learning algorithms. *PLOS ONE*, 16(10), e0258125. doi:10.1371/journal.pone.0258125
- Gulliver, J.S. 2007. Introduction to chemical transport in the environment. Cambridge University Press.
- Hasanpour Kashani, M., Ghorbani, M.A., Dinpashoh, Y. et al. 2014. Comparison of Volterra Model and Artificial Neural Networks for Rainfall–Runoff Simulation. *Nat Resour Res* 23, 341–354. doi:10.1007/s11053-014-9235-y
- Hoffmann, T., Thorndycraft, V.R., Brown, A.G., Coulthard, T.J., Damnati, B., Kale, V.S., Middelkoop, H., Notebaert, B., and Walling, D.E. 2023. Pristine levels of suspended sediment in a large river. *Earth Surface Dynamics*, 11, 287-301. doi:10.5194/esurf-11-287-2023
- Bezak, N., Lebar, K., Bai, Y., and Rusjan, S. 2025. Using machine learning to predict suspended sediment transport under climate change. *Water Resources Management*, 39, 3311-3326. doi:10.1007/s11269-025-04108-7
- Biazar, S. M., Fard, A.F., Singh, V.P. et al. 2020. Estimation of Evaporation from Saline-Water with More Efficient Input Variables. *Pure Appl. Geophys.* 177, 5599–5619. doi:10.1007/s00024-020-02570-5
- Collins, A. L., Zhang, Y., Walling, D. E., Grenfell, S. E., and Smith, P. 2010. Tracing sediment loss from eroding farm tracks using a geochemical fingerprinting procedure combining local and genetic algorithm optimisation. *Science of the Total Environment*, 408(22), 5461-5471. doi: 10.1016/j.scitotenv.2010.07.066
- Darabi, H., Mohamadi, S., Karimidastenaie, Z., Kisi, O., and Ehteram, M. 2021. Prediction of daily suspended sediment load (SSL) using new optimization algorithms and soft computing models. *Soft Computing*, 25, 7609-7626
- Dethier, E.N., Magilligan, F.J., Renshaw, C.E., and Nislow, K.H. 2022. Rapid changes to global river suspended sediment flux by human activities. *Science*, 376(6595), 1447-1451. doi:10.1126/science.abn7980
- Dinpazhoh, Y., Sattari, M., E, S., and Darbandi, S. 2017. 'Optimum Operation of Reservoir Using the Genetic Algorithm and Particle Swarm Optimization (Case Study: Alavian Dam)', *Water and Soil Science*, 27(2), pp. 17-29.
- Ebtehaj, I., and Bonakdari, H. 2017. Design of a fuzzy differential evolution algorithm to predict non-deposition sediment transport. *Applied Water Science*, 7, 4287-4299. doi:10.1007/s13201-017-0562-0
- Emami, S., and Choopan, Y. 2019. An application of imperialist competitive algorithm in estimating sediment suspended load. *Hydrogeology*, 4(1), 70-79. doi:10.22034/hydro.2019.7310 [In Persian]
- Emami, S., Choopan, Y., and Parsa, J. 2018. Modeling the groundwater level of the Miandoab plain using artificial neural network method and election and genetic algorithms. *Journal of Ecohydrology*, 5(4), 1175-1189. doi:10.22059/ije.2018.258644.889 [In Persian]
- Emami, S.K., Saadat, A.M., and Hamidifar, H. 2024. Microplastic pollution: Analytical techniques, policy landscape, and integrated strategies for sustainable environmental stewardship. In S.C. Izah, M.C. Ogwu, A. Loukas, and H. Hamidifar (Eds.), *Water crises and sustainable management in the Global South* (pp. 237-270). Springer. doi:10.1007/978-981-97-4966-9_11

- Li, S., and Yang, J. 2022. Modelling of suspended sediment load by Bayesian optimized machine learning methods with seasonal adjustment. *Engineering Applications of Computational Fluid Mechanics*, 16(1), 1883-1901. doi:10.1080/19942060.2022.2121944
- Melesse, A.M., Ahmad, S., McClain, M.E., Wang, X., and Lim, Y.H. 2011. Suspended sediment load prediction of river systems: An artificial neural network approach. *Agricultural Water Management*, 98(5), 855-866. doi:10.1016/j.agwat.2010.12.012
- Meshram, S.G., Singh, V.P., Kisi, O., Karimi, V., and Meshram, C. 2020. Application of artificial neural networks, support vector machine and multiple model-ANN to sediment yield prediction. *Water Resources Management*, 34(15), 4561-4575. doi:10.1007/s11269-020-02672-8
- Moradinejad, A. 2024. Modeling estimation of river suspended sediment amount using support vector regression and data control group method. *Water and Soil Science*, 34(2), 203-222. doi:10.22034/ws.2023.57435.2526 [In Persian]
- Naseri F, Azari M, and Dasturani M T. 2019. Optimizing coefficients of Sediment rating curve equation using Genetic Algorithm (Case study: Ghazaghli and Bagh abbasi stations), *Irrigation and Water Engineering*, 9(3), pp. 82-98. doi: 10.22125/iwe.2019.88672 [In Persian]
- Nda, M., Shabako, J.G., Yakubu, M., Godwin, A., and Justice, A.N. 2023. Riverbank erosion, sediment transport and reservoir sedimentation: An overview. *Savannah Journal of Science and Engineering Technology*, 1(4), 210-216.
- Nikoo, M., Razavi, S. A., and Hadzima-Nyarko, M. 2018. Artificial Neural Network Combined with Imperialist Competitive Algorithm for Determination of River Sediments. *Fresenius Environ Bull*, 27(7), 4658-4667.
- Noori, R., Karbassi, A.R., Ashrafi, K., Farokhnia, A., and Sabahi, M.S. 2012. Active and online prediction of BOD5 in river systems using reduced-order support vector machine. *Environmental Earth Sciences*, 67, 141-149. doi:10.1007/s12665-011-1487-9
- Parsons, A.J., Cooper, J., and Wainwright, J. 2015. What is suspended sediment? *Earth Surface Processes and Landforms*, 40(10), 1417-1420. doi:10.1002/esp.3730
- Razavizadeh, S., and Dargahian, F. 2018. Optimization of artificial neural network structure in prediction of sediment discharge using Taguchi method. *Journal of Water Management Science and Engineering*, 12(43), 89-97. doi:10.1001.1.20089554.1397.12.43.10.3 [In Persian]
- Joshi, B., Singh, V.K., Vishwakarma, D.K., Ghorbani, M.A., Kim, S., Gupta, S., Chandola, V.K., Rajput, J., Chung, I.M., Yadav, K.K., Mirzania, E., Al-Ansari, N., and Mattar, M.A. 2024. A comparative survey between cascade correlation neural network (CCNN) and feedforward neural network (FFNN) machine learning models for forecasting suspended sediment concentration. *Scientific Reports*, 14, 10638. doi:10.1038/s41598-024-61339-1
- Jothiprakash, V., and Garg, V. 2009. Reservoir sedimentation estimation using artificial neural network. *Journal of Hydrologic Engineering*, 14(9), 1035-1040. doi:10.1061/(ASCE)HE.1943-5584.0000075
- Khanduzi, R., Ebrahimzadeh, A., and Ebrahimzadeh, Z. 2024. A combined Bernoulli collocation method and imperialist competitive algorithm for optimal control of sediment in the dam reservoirs. *AUT Journal of Mathematics and Computing*, 5(1), 71-80. doi:10.22060/ajmc.2023.21635.1109
- Kisi, O. 2007. Evapotranspiration modelling from climatic data using a neural computing technique. *Hydrological Processes*, 21(14), 1925-1934. doi:10.1002/hyp.6403
- Kisi, O., and Ozkan, C. 2017. A new approach for modeling sediment-discharge relationship: Local weighted linear regression. *Water Resources Management*, 31(1), 1-23. doi:10.1007/s11269-016-1481-9
- Kisi, O., and Shiri, J. 2012. River suspended sediment estimation by climatic variables implication: Comparative study among soft computing techniques. *Computers & Geosciences*, 43, 73-82. doi:10.1016/j.cageo.2012.02.007
- Kisi, O., Heddam, S., Parmar, K.S., Malik, A., and Karimi, B. 2024. Improved monthly streamflow prediction using integrated multivariate adaptive regression spline with K-means clustering: Implementation of reanalyzed remote sensing data. *Stochastic Environmental Research and Risk Assessment*, 38, 2489-2519. doi:10.1007/s00477-024-02692-5
- Knighton, D. 2014. *Fluvial forms and processes: A new perspective*. Routledge.
- Latif, S. D., Chong, K. L., Ahmed, A. N., Huang, Y. F., Sherif, M., and El-Shafie, A. 2023. Sediment load prediction in Johor river: deep learning versus machine learning models. *Applied Water Science*, 13(3), 79. doi: 10.1007/s13201-023-01874-w
- Le, X.H., Ho, H.V., Lee, G., and Jung, S. 2024. Quantifying predictive uncertainty and feature selection in machine learning river bed load predictions. *Water*, 16(14), 1945. doi:10.3390/w16141945

- Shahi Nejad, B., and Dehghani, R. 2017. Evaluation of support vector machine model performance in estimating suspended sediment. *Irrigation and Water Engineering*, 8(1), 30–42. [In Persian]
- Sharafati, A., Haji Seyed Asadollah, S. B., Motta, D., and Yaseen, Z. M. 2020. Application of newly developed ensemble machine learning models for daily suspended sediment load prediction and related uncertainty analysis. *Hydrological Sciences Journal*, 65(12), 2022–2042. doi:10.1080/02626667.2020.1786571
- Smits, A.J.M., Nienhuis, P.H., and Leuven, R.S.E.W. 2000. New approaches to river management. Backhuys Publishers.
- Syed, A.U., Konieczny, J.T., and Wright, S.A. 2005. A pre-dam-removal assessment of sediment transport for four dams on the Klamath River using sediment transport modeling. U.S. Geological Survey Scientific Investigations Report 2004-5178.
- Teixeira, L.C., Mariani, P.P., Pedrollo, O.C., and Avila, L.F. 2020. Artificial neural network and fuzzy inference system models for forecasting suspended sediment and turbidity in basins at different scales. *Water Resources Management*, 34, 3709-3723. doi:10.1007/s11269-020-02647-9
- Vanoni, V.A. (Ed.). 2006. Sedimentation engineering. American Society of Civil Engineers. doi:10.1061/9780784408230
- Vázquez-Tarrio, D., Ruiz-Villanueva, V., Garrote, J., Benito, G., Calle, M., Lucía, A., and Díez-Herrero, A. 2024. Effects of sediment transport on flood hazards: Lessons learned and remaining challenges. *Geomorphology*, 446, 108976. doi:10.1016/j.geomorph.2023.108976
- Yadav, A., Chatterjee, S., and Equeenuddin, S.M. 2018. Suspended sediment yield estimation using genetic algorithm-based artificial intelligence models: Case study of Mahanadi River, India. *Hydrological Sciences Journal*, 63(8), 1162-1182. doi:10.1080/02626667.2018.1483581
- Yadav, A., Hasan, M. K., Joshi, D., Kumar, V., Aman, A. H. M., Alhumyani, H., Alzaidi, M. S., and Mishra, H. 2022. Optimized Scenario for Estimating Suspended Sediment Yield Using an Artificial Neural Network Coupled with a Genetic Algorithm. *Water*, 14(18), 2815. doi:10.3390/w14182815
- Zangeneh Asadi, M.A., Goli Mokhtari, L., Zandi, R., and Naemitabar, M. 2025. Modeling, evaluation and forecasting of suspended sediment load in Kal-e Shur River, Sabzevar Basin, in northeast of Iran. *Applied Water Science*, 15, 44. doi:10.1007/s13201-025-02361-0
- Zhou, Y., Wang, H., and Yan, Y. 2023. Sediment transport trend and its influencing factors in a large
- Saadat, A.M., and Mazaheri, M. 2022. Backward solution (in-time) of the pollution transport equation in river using group preserving scheme. *Journal of Ferdowsi Civil Engineering*, 35(4), 35–52. doi:10.22067/jfcei.2022.77645.1165 [In Persian]
- Saadat, A.M., and Mazaheri, M. 2023. Application of one-step group preserving scheme in various contaminant transport modeling in rivers. *Iranian Journal of Soil and Water Research*, 54(11), 1627–1646. doi:10.22059/ijswr.2023.365226.669572 [In Persian]
- Saadat, A.M., and Mazaheri, M. 2024a. Forward and inverse river contaminant transport modeling using group preserving scheme. *Physics of Fluids*, 36(9), 094102. doi:10.1063/5.0221578
- Saadat, A.M., and Mazaheri, M. 2024b. Using one-step group preserving schemes for contaminant transport modeling in rivers. *Iranian Journal of Soil and Water Research*, 54(11), 1627-1646. doi:10.22059/ijswr.2023.365226.669572 [In Persian]
- Saadat, A.M., Emami, S.K., and Hamidifar, H. 2024. A review on storage process models for improving water quality modeling in rivers. *Hydrology*, 11(11), 187. doi:10.3390/hydrology11110187
- Saadat, A.M., Emami, S.K., and Mazaheri, M. 2025b. Evaluating the impact of storage zones on backward contaminant transport: A comparative study of the classic advection-dispersion equation and storage zone models in riverine systems. *Science of The Total Environment*, 973, 179176. doi:10.1016/j.scitotenv.2025.179176
- Saadat, A.M., Emami, S.K., Hamidifar, H., and Fardadi Shilsar, M.J. 2025a. Microplastic pollution in water systems of the Global South: A review. *Journal of Contaminant Hydrology*, 267, 104729. doi:10.1016/j.jconhyd.2025.104729
- Saadat, A.M., Mazaheri, M., and Mohammad Vali Samani, J. 2022. Backward (in-time) inverse solution of river contaminant transport equation using the group preserving scheme. *Journal of Ferdowsi Civil Engineering*, 35(4), 35–52. doi:10.22067/jfcei.2022.77645.1165 [In Persian]
- Sanikhani, H., Kisi, O., Nikpour, M.R. et al. 2012. Estimation of Daily Pan Evaporation Using Two Different Adaptive Neuro-Fuzzy Computing Techniques. *Water Resour Manage* 26, 4347–4365. doi:10.1007/s11269-012-0148-4
- Seyedian, S.M., Rohani, H., Fathabadi, A., and Javadi Alinezhad, M. 2017. Stochastic modeling of sediment yield using random forest and quantile regression. *Journal of Water and Soil Conservation*, 24(4), 103–122. doi:10.22069/jwsc.2017.12600.2732 [In Persian]

