

Assessment of the performance of data-driven models in estimating suspended sediment in high-sediment rivers: a case study of the Hootan gauging station, Atrak River, Maraveh Tappeh study unit

Mahmoudreza Tabatabaei^{1*} and Mohammadreza Gharib Reza¹

¹ Associate Professor, Soil Conservation and Watershed Management Research Institute, Agricultural Research, Education and Extension Organization (AREEO), Tehran, Iran

Received: 23 June 2025

Accepted: 04 October 2025

Extended abstract

Introduction

Accurate estimation of suspended sediment in high-sediment rivers, especially in semi-arid regions such as the Golestan River watershed, poses significant challenges in water resource management and sediment control in reservoir dams. The increase in sediment concentration not only affects water quality but also leads to considerable economic and environmental damage by reducing the lifespan of hydraulic structures and altering river morphology. In this context, the Atrak River, as one of the most important sediment sources in northeastern Iran, serves as a prominent example of these challenges. Given the limitations of traditional direct measurement methods and labor-intensive physical models, the development of high-accuracy data-driven models emerges as an efficient solution for monitoring and predicting sediment.

Materials and methods

In this study, to estimate suspended sediment in the Atrak River at the Hootan gauging station, a combination of classical and intelligent methods was employed, including: (1) sediment rating curves (using both the midpoint and linear methods), (2) neural networks (MLP and SOM), (3) deep learning models, and (4) ensemble learning methods. The modeling process was carried out in three main stages: first, using the Random Forest algorithm, the key variables affecting sediment (including flow rate, daily precipitation, and their lagged values) were identified. Then, the data were divided into homogeneous groups through clustering, allowing for balanced sampling from each cluster to create homogeneous training (70%), evaluation (15%), and testing (15%) datasets. To enhance the model's efficiency, various strategies were employed, including optimizing objective functions and managing data skewness. After a comprehensive evaluation of the models using standard criteria, the ensemble learning model XGBoost was selected as the best model, which was ultimately used to reconstruct and complete suspended sediment data for a 40-year period (1982-2021).

Results and discussion

The comparison of the results of various models in this study indicated that the ensemble learning model (XGBoost) was recognized as the selected model among other data-driven models for estimating suspended sediment in the Atrak River, having the lowest error rates (MAE of 10,652 tons per day and RMSE of 36,219 tons per day) and the highest performance index (NSE of 0.87). This model not only outperformed other machine learning methods (with the MLP neural network achieving NSE=0.80 and the deep learning model achieving NSE=0.84), but it also demonstrated significant superiority over classical methods such as the midpoint sediment rating curve (with MAE of 20,909 tons per day and RMSE of 45,632 tons per day). A detailed analysis of the results indicates that this superiority is primarily due to the ability of the XGBoost model to manage highly skewed data and identify complex nonlinear relationships between hydrological variables. In this regard, the model's bias (PBias) decreased significantly by approximately 15.5% from -18.93 in the sediment rating curve model to -3.52 in the XGBoost model, supporting this point. On the other hand, a comparative analysis with similar studies in other high-sediment rivers shows that the combined approach used in this research (utilizing clustering before model training and optimizing objective functions) has led to improved predictions. However, it is important to note that the effectiveness of these models largely depends on the quality and completeness of the input data, and long-term hydrological changes may necessitate periodic recalibration (training) of the models. These limitations

* Corresponding author: taba1345@hotmail.com

provide a foundation for future research aimed at developing models that are more adaptable to environmental changes.

Conclusion

This research examined the performance of data-driven models in estimating suspended sediment in the Atrak River at the Hootan gauging station, yielding significant results. The ensemble learning model XGBoost was identified as the best option, demonstrating high accuracy and minimal error in predicting changes in suspended sediment. These findings highlight the importance of utilizing advanced machine learning techniques in analyzing hydrological data and emphasize the need for a connection between data science and water resource management. Key points in this study include the high skewness of suspended sediment data and the specific hydrological conditions of the Atrak River, which arise from gully erosion and other natural phenomena in the region. This skewness and the sharp fluctuations in sediment concentration pose major challenges for accurate modeling. However, the combined approach used in this study provided a significant improvement in prediction accuracy compared to traditional and classical methods. This indicates that employing clustering and data optimization can serve as effective tools in future research. Nonetheless, the results of this study should be interpreted with caution, as the quality of data and environmental changes may impact model performance. Therefore, this research not only aids in a better understanding of sedimentation processes but also lays the groundwork for future studies aimed at enhancing and developing prediction models that are more adaptable to changing environmental conditions. In the end, the results of this research can serve as a model for simulating suspended sediment in other hydrometric stations across the country, providing valuable insights for researchers and experts in this field.

Keywords: Data-driven models, Ensemble learning, Gradient boosting, Machine learning methods, Prediction, XGBoost

Cite this article: Tabatabaei, M., Gharib Reza, M., 2026. Assessment of the performance of data-driven models in estimating suspended sediment in high-sediment rivers: a case study of the Hootan gauging station, Atrak River, Maraveh Tappeh study unit. *Water. Eng. Manag.*18(3), 297-320.

© 2026, The Author(s). Published by Soil Conservation and Watershed Management Research Institute (SCWMRI). This is an open-access article distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0>)



ارزیابی کارایی مدل‌های داده مبنا در برآورد رسوب معلق رودخانه‌های پر رسوب، مطالعه موردی ایستگاه آب‌سنجی هوتن رودخانه اترک، واحد مطالعاتی مراوه تپه

محمودرضا طباطبائی^{۱*} و محمدرضا غریب رضا^۱

^۱ دانشیار پژوهشی پژوهشکده حفاظت خاک و آبخیزداری، سازمان تحقیقات، آموزش و ترویج کشاورزی، تهران، ایران

تاریخ پذیرش: ۱۴۰۴/۰۷/۱۲

تاریخ دریافت: ۱۴۰۴/۰۴/۰۲

چکیده مبسوط

مقدمه

برآورد دقیق رسوب معلق در رودخانه‌های پررسوب، به‌ویژه در مناطق نیمه‌خشک مانند حوزه گرگانرود، از چالش‌های اساسی در مدیریت منابع آب و کنترل رسوبات مخازن سدها محسوب می‌شود. افزایش غلظت رسوبات نه تنها بر کیفیت آب تأثیر می‌گذارد، بلکه با کاهش عمر مفید سازه‌های هیدرولیکی و تغییر مورفولوژی رودخانه، خسارات اقتصادی و زیست‌محیطی قابل توجهی به دنبال دارد. در این میان، رودخانه اترک به‌عنوان یکی از مهم‌ترین منابع رسوبی در شمال شرق ایران، نمونه‌ای بارز از این چالش‌هاست. با توجه به محدودیت‌های روش‌های سنتی اندازه‌گیری مستقیم و مدل‌های فیزیکی پرزحمت، توسعه مدل‌های داده‌مبنا با دقت بالا به‌عنوان راهکاری کارآمد برای پایش و پیش‌بینی رسوب مطرح می‌شود.

مواد و روش‌ها

در این پژوهش، برای برآورد رسوب معلق رودخانه اترک در ایستگاه هوتن، از ترکیبی از روش‌های کلاسیک و هوشمند شامل: (۱) منحنی سنج رسوب (با دو روش حد وسط دسته‌ها و یک خطی)، (۲) شبکه‌های عصبی (MLP و SOM)، (۳) مدل‌های یادگیری عمیق و (۴) روش‌های یادگیری جمعی استفاده شد. فرایند مدل‌سازی در سه مرحله اصلی انجام پذیرفت. ابتدا با بهره‌گیری از الگوریتم جنگل تصادفی، متغیرهای کلیدی مؤثر بر رسوب (شامل دبی جریان، بارش روزانه و مقادیر تأخیری آنها) شناسایی شدند. سپس داده‌ها با روش خوشه‌بندی به گروه‌های همگن تقسیم شدند که این امر امکان نمونه‌برداری متوازن از هر خوشه برای ایجاد مجموعه‌های همگن آموزش (۷۰ درصد)، ارزیابی (۱۵ درصد) و آزمون (۱۵ درصد) را فراهم نمود. به‌منظور افزایش کارایی مدل‌ها، از راهکارهای مختلفی شامل بهینه‌سازی توابع هدف و مدیریت چولگی داده‌ها استفاده شد. پس از ارزیابی جامع مدل‌ها با معیارهای استاندارد، مدل یادگیری جمعی XGBoost به‌عنوان مدل برتر انتخاب شد که در نهایت برای بازسازی و تکمیل داده‌های رسوب معلق در دوره ۴۰ ساله (۱۴۰۱-۱۳۶۰) به کار گرفته شد.

نتایج و بحث

مقایسه نتایج مدل‌های مختلف در این پژوهش نشان داد که مدل یادگیری جمعی (XGBoost) با داشتن کمترین میزان خطا (MAE برابر با ۱۰۶۵۲ تن در روز و RMSE برابر ۳۶۲۱۹ تن در روز) و بیشترین شاخص کارایی (NSE برابر با ۰/۸۷)

به عنوان مدل منتخب از بین سایر مدل‌های داده مبنای برای برآورد رسوب معلق در رودخانه اترک شناخته شد. این مدل نه تنها در مقایسه با سایر روش‌های یادگیری ماشین مانند شبکه عصبی MLP (با $NSE=0.80$) و مدل یادگیری عمیق با $NSE=0.84$ عملکرد بهتری داشت، بلکه نسبت به روش‌های کلاسیک مانند منحنی سنجی حد وسط دسته‌ها با MAE برابر با ۲۰۹۰۹ تن در روز و RMSE برابر ۴۵۶۳۲ تن در روز، نیز برتری قابل توجهی نشان داد. تحلیل دقیق نتایج حاکی از آن است که این برتری عمدتاً ناشی از توانایی مدل XGBoost در مدیریت داده‌های با چولگی بالا و شناسایی رابطه‌های غیرخطی پیچیده بین متغیرهای هیدرولوژیکی است. در این رابطه مقدار بایاس (PBias) مدل با کاهش تقریباً ۱۵/۵ درصدی از ۱۸/۹۳- در مدل منحنی سنجی رسوب به ۳/۵۲- در مدل XGBoost مؤید این نکته است. از سوی دیگر، مقایسه تطبیقی با مطالعات مشابه نشان می‌دهد که رویکرد ترکیبی به کار گرفته شده در این پژوهش (استفاده از خوشه‌بندی پیش از آموزش مدل و بهینه‌سازی توابع هدف) منجر به بهبود پیش‌بینی شده است. با این حال، باید توجه داشت که کارایی این مدل‌ها تا حد زیادی به کیفیت و کامل بودن داده‌های ورودی وابسته است و تغییرات هیدرولوژیکی بلندمدت ممکن است نیاز به بازواسنجی (آموزش) دوره‌ای مدل‌ها را ایجاد کند. این محدودیت‌ها زمینه‌ای برای تحقیقات آتی در جهت توسعه مدل‌های سازگارتر با تغییرات محیطی فراهم می‌سازد.

نتیجه‌گیری

این پژوهش به بررسی کارایی مدل‌های داده‌مبنای در برآورد رسوب معلق در رودخانه اترک، ایستگاه آب‌سنجی هوتن پرداخته و به نتایج قابل توجهی دست یافته است. مدل یادگیری جمعی XGBoost به عنوان بهترین گزینه شناسایی شد که با دقت بالا و کمترین خطا، توانایی پیش‌بینی تغییرات رسوب معلق را به نمایش گذاشت. این یافته‌ها نشان‌دهنده اهمیت استفاده از تکنیک‌های پیشرفته یادگیری ماشین در تحلیل داده‌های هیدرولوژیکی است و بر لزوم پیوند بین علم داده و مدیریت منابع آب تأکید می‌کند. از نکات کلیدی در این تحقیق، چولگی بالای داده‌های رسوب معلق و شرایط خاص هیدرولوژیکی رودخانه اترک است که ناشی از فرسایش‌های خندقی و پدیده‌های طبیعی دیگر در این منطقه است. این چولگی و تغییرات شدید در غلظت رسوب، چالش‌های عمده‌ای را برای مدلسازی دقیق به وجود می‌آورد. با این حال، رویکرد ترکیبی به کار گرفته در این مطالعه، بهبود قابل توجهی در دقت پیش‌بینی نسبت به روش‌های سنتی و کلاسیک فراهم آورد. این امر حاکی از آن است که استفاده از خوشه‌بندی و بهینه‌سازی داده‌ها می‌تواند به عنوان ابزاری مؤثر در تحقیقات آینده مورد توجه قرار گیرد. با این حال، نتایج این تحقیق باید با احتیاط تفسیر شوند. چرا که کیفیت داده‌ها و تغییرات محیطی ممکن است بر عملکرد مدل‌ها تأثیر بگذارد. بنابراین، این مطالعه نه تنها به درک بهتر فرایندهای رسوبگذاری کمک می‌کند، بلکه زمینه‌ساز تحقیقات آتی در جهت بهبود و توسعه مدل‌های پیش‌بینی سازگارتر با شرایط متغیر محیطی خواهد بود. در نهایت نتایج این تحقیق می‌تواند به عنوان یک کار الگویی در شبیه‌سازی رسوب معلق در دیگر ایستگاه‌های هیدرومتری کشور مورد استفاده محققین و کارشناسان این حوزه قرار گیرد.

واژه‌های کلیدی: پیش‌بینی، تقویت گرادیان، رسوب معلق، گرادیان بوستینگ، مدل‌های داده‌مبنای، مدل‌های یادگیری ماشین، یادگیری جمعی

مقدمه

تغییرات، ناشی از عوامل طبیعی مانند فرسایش خندقی، بارندگی‌های شدید و فعالیت‌های انسانی است که دقت مدل‌های سنتی برآورد رسوب را تحت تأثیر قرار می‌دهد. در ادامه، نتایج برخی از پژوهش‌های اخیر در زمینه شبیه‌سازی و برآورد رسوب معلق رودخانه‌ها ذکر شده است.

برآورد دقیق رسوب معلق در رودخانه‌های پرسوب از چالش‌های اساسی در مدیریت منابع آب و کنترل فرسایش است. رودخانه اترک، به ویژه در محدوده ایستگاه آب‌سنجی هوتن (واحد مطالعاتی مراوه‌تپه)، به دلیل شرایط هیدرولوژیکی خاص و فرسایش‌پذیری بالا، با نوسانات شدید غلظت رسوب مواجه است. این

مدل با تابع فعال‌سازی Huber و تابع زیان ReLU بهترین عملکرد را داشت. این تحقیق تأکید می‌کند که استفاده از داده‌های بارش ماهواره‌ای می‌تواند کارایی مدل‌های پیش‌بینی رسوب را بهبود بخشد و پیشنهاد می‌کند که تجربیات حاصل در دیگر ایستگاه‌های کشور نیز به‌کار گرفته شود.

در پژوهشی که توسط Üneş et al., (2025) در رودخانه West Branch Neversink در ایالات متحده انجام شد، پیش‌بینی رسوب معلق با استفاده از هوش مصنوعی بررسی شد تا دقت پیش‌بینی در شرایط آب و هوایی متغیر افزایش یابد. یکی از مهم‌ترین دستاوردهای این تحقیق، عملکرد برتر در پیش‌بینی روزانه با استفاده از داده‌های سال‌های ۲۰۲۱ تا ۲۰۲۳ بود که این نتایج برای نظارت بر سیلاب‌ها بسیار مفید به نظر می‌رسد. مدل‌های مورد استفاده شامل شبکه‌های عصبی مصنوعی^۳، رگرسیون خطی چندگانه^۴ و درخت تصمیم M5 بودند. مقایسه‌ها نشان دادند که شبکه‌های عصبی مصنوعی (ANN) بهترین عملکرد را در مقایسه با رگرسیون خطی چندگانه (MLR) و درخت تصمیم M5 داشته و خطای کمتری را ارائه می‌دهند. همچنین، داده‌های به‌کاررفته شامل ورودی‌هایی نظیر دمای هوا، بارش، جریان و کدورت بودند که به بهبود دقت پیش‌بینی کمک کردند.

در پژوهشی که توسط Doroudi and Sharafati, (2021) در حوزه آبخیز سد کوثر واقع در جنوب غربی ایران انجام شد، از روش شبکه عصبی مصنوعی (ANN) برای پیش‌بینی بار رسوب معلق (SSL) استفاده شد. در این مطالعه، دبی رودخانه و بارش به‌عنوان ویژگی‌های ورودی در نظر گرفته شدند و مدل‌های پیش‌بینی ANN-OTLBO و ANN-PSO با استفاده از الگوریتم‌های فراابتکاری OTLBO و PSO توسعه یافتند. نتایج نشان داد که مدل ANN-OTLBO بهترین عملکرد را در پیش‌بینی بار رسوب معلق داشت و ترکیب الگوریتم‌های فراابتکاری با شبکه عصبی مصنوعی به‌طور معنی‌داری دقت پیش‌بینی را افزایش داد.

در پژوهشی که توسط Shadkani et al., (2025) در رودخانه Mississippi و در ایستگاه‌های Chester و

در پژوهشی که توسط Taşar et al., (2024) در رودخانه Cuyahoga واقع در اوهایو، ایالات متحده انجام شد، به بررسی پیش‌بینی روزانه رسوب معلق با استفاده از روش‌های یادگیری ماشین پرداخته شد تا دقت برآورد در رودخانه‌های پررسوب بهبود یابد. از مهم‌ترین دستاوردهای این تحقیق، ارتقاء قابل توجه دقت پیش‌بینی با بهره‌گیری از مدل‌های داده‌محور بود که بر اساس داده‌های تاریخی بین سال‌های ۱۹۶۹ تا ۱۹۷۴ ارزیابی شد. مدل‌های مورد استفاده شامل ماشین‌های بردار پشتیبان (SVM)، درخت تصمیم M5 (M5T) و رگرسیون خطی چندگانه (MLR) بودند. مقایسه‌ها نشان داد که SVM و M5T عملکرد بهتری نسبت به MLR دارند، به‌طوری‌که RMSE کمتری را ارائه داده و R^2 بالاتری را به نمایش گذاشتند. همچنین، ورودی‌های اصلی شامل جریان رودخانه، غلظت رسوب و دمای آب بودند که این مدل‌ها را برای کاربردهای هیدرولوژیکی مناسب‌تر کرد.

در پژوهشی که توسط Jarbais and Harshavardhanan, (2025) در رودخانه Ganga در بنارس، هند انجام شد، پیش‌بینی غلظت رسوب معلق روزانه با استفاده از مدل‌های یادگیری ماشین مقایسه شد تا بهترین ابزار برای نظارت محیطی در رودخانه‌های پررسوب شناسایی شود. یکی از مهم‌ترین دستاوردهای این تحقیق، بهبود دقت پیش‌بینی با تمرکز بر داده‌های روزانه یک ساله بود که اهمیت مدل‌های هوشمند در مدیریت آبخیز را به‌خوبی نمایان ساخت. مدل‌های مورد بررسی شامل شبکه عصبی پس انتشار روبجلو^۱ و ماشین بردار پشتیبان^۲ بودند. مقایسه‌ها نشان داد که FFNN عملکرد بهتری نسبت به SVM دارد، به‌طوری‌که با RMSE پایین‌تر، ضریب همبستگی R بالاتر و NSE بهتر همراه بود.

در پژوهشی دیگر، Tabatabaei et al., (2023) به بررسی تأثیر بارش روزانه به‌دست‌آمده از ماهواره CHIRPS بر شبیه‌سازی رسوب معلق رودخانه قره‌چای پرداختند. در این پژوهش با استفاده از شبکه عصبی مصنوعی پرسپترون چند لایه و متغیرهای دبی جریان و بارش، نه مدل مختلف ایجاد شد. نتایج نشان داد که

³ Artificial neural networks (ANNs)

⁴ Multiple linear regression (MLR)

¹ Feed-Forward Backpropagation (FFNN)

² Support Vector Machine (SVM)

مدل‌ها را در مدیریت آبخیز تأیید کرد. مدل‌های مورد استفاده شامل ANFIS و ANN بودند؛ مقایسه‌ها نشان داد که ANFIS عملکرد بهتری نسبت به ANN دارد، به‌طوری‌که خطای کمتری و همبستگی بالاتری را به نمایش گذاشت.

در پژوهشی که توسط Ghanbari-Adivi (2025) در رودخانه تالار در مازندران، ایران انجام شد، یک مدل هیبریدی جدید برای پیش‌بینی بار رسوب معلق معرفی شد تا چالش‌های رودخانه‌های پرسوب در شرایط آب و هوایی متغیر حل شود. یکی از مهم‌ترین دستاوردهای این تحقیق، دستیابی به دقت بالا با MAE پایین (۵۲۵)، NSE بالا (۰/۹۸) و PBIAS کم (۴) بود که این مدل را برای مدیریت سیل مناسب ساخت. ابزارهای به‌کاررفته شامل مدل هیبریدی ANFIS⁴-M5T و ورودی‌های تأخیری بارش، دبی و رسوب بودند. هرچند این مدل به‌طور مستقیم با دیگر مدل‌ها مقایسه نشد، اما نشان داد که ترکیب فازی و تصمیم‌گیری درختی می‌تواند دقت پیش‌بینی را نسبت به مدل‌های تک به‌طور قابل توجهی بهبود بخشد.

در پژوهش که توسط Mohamadi (2019) در بررسی ایستگاه‌های هیدرومتری بالادست رودخانه هلیل‌رود انجام شد، از مدل‌های مختلف شامل مدل‌های منحنی سنج رسوب یک خطی و دو خطی، روش حد وسط دسته‌ها، و همچنین مدل‌های جعبه سیاه نظیر شبکه عصبی مصنوعی و سامانه استنتاج عصبی-فازی برای برآورد میزان رسوب معلق استفاده شد. نتایج نشان داد که مدل‌های عصبی-فازی در ایستگاه‌های پل بافت، هنجان و کناروئیه بهترین عملکرد را داشتند. همچنین، مدل شبکه عصبی مصنوعی تابع پایه شعاعی در ایستگاه میدان و روش منحنی سنج رسوب دو خطی در ایستگاه چشمه عروس به‌عنوان بهترین مدل‌ها برای شبیه‌سازی میزان رسوب معلق شناخته شدند. این نتایج نشان‌دهنده قابلیت بالای این مدل‌ها در تخمین بار رسوبی با استفاده از داده‌های دبی آب است.

در پژوهشی که توسط and Bahnam, (2023) در حوضه رودخانه Greater Zab در عراق انجام شد، پیش‌بینی بار رسوب معلق ماهانه با استفاده

Thebes انجام شد، پیش‌بینی غلظت رسوب معلق با استفاده از یک مدل هیبریدی جدید بررسی شد تا دقت پیش‌بینی در جریان‌های رودخانه‌ای افزایش یابد. یکی از مهم‌ترین دستاوردهای این تحقیق، دستیابی به ضریب همبستگی بالا (۰/۹۷۶ و ۰/۹۶۰) با داده‌های USGS در بازه زمانی ۲۰۰۷ تا ۲۰۱۷ بود. مدل‌های مورد استفاده شامل FFNN، مدل شبکه عصبی حافظه طولانی کوتاه مدت^۱، گرادیان کاهشی تصادفی^۲، شبکه عصبی شعاعی پایه^۳ و مدل هیبریدی SFFNN-LSTM بودند. مقایسه‌ها نشان داد که مدل SFFNN-LSTM-6 عملکرد بهتری نسبت به سایر مدل‌ها دارد و خطای کمتری را ارائه می‌دهد. این رویکرد همچنین به‌عنوان ابزاری مناسب برای شبیه‌سازی هیدرولوژیکی پیشنهاد شده است.

در پژوهشی که توسط Bari et al., (2024) به عنوان بررسی جهانی روندهای تخمین رسوب معلق از سال ۲۰۱۱ تا ۲۰۲۲ انجام شد، مدل‌های یادگیری ماشین برای برآورد در رودخانه‌های پرسوب ارزیابی شدند تا چالش‌های هیدرولوژیکی جهانی شناسایی شود. یکی از مهم‌ترین دستاوردهای این تحقیق، تأکید بر عملکرد برتر مدل‌هایی مانند LSTM و ANFIS نسبت به مدل‌های سنتی بود که محدودیت‌های فرکانس نمونه‌برداری را حل کردند. مدل‌های مورد بررسی شامل SWAT، WEPP، LSTM و ANFIS بودند. مقایسه‌ها نشان داد که مدل‌های یادگیری ماشین دقت بالاتری نسبت به مدل‌های فیزیکی مانند SWAT دارند. همچنین، ابزارهای به‌کاررفته شامل بررسی محدودیت‌ها برای بهبود شبیه‌سازی بودند که به درک بهتر عملکرد این مدل‌ها کمک کرد.

در پژوهشی که توسط Mousavi et al., (2024) در رودخانه دز در خوزستان، ایران انجام شد، تخمین بار رسوب معلق با استفاده از مدل‌های هوشمند داده‌محور مورد بررسی قرار گرفت تا شبیه‌سازی دقیق‌تری برای رودخانه‌های پرسوب ارائه شود. یکی از مهم‌ترین دستاوردهای این تحقیق، بهبود دقت تخمین با تمرکز بر معیارهای ارزیابی مانند R²، RMSE و Nash-Sutcliffe Efficiency (NS) بود که کاربرد عملی

³ Radial Basis Function (RBF)

⁴ Adaptive Neuro-Fuzzy Inference System

¹ Long Short Term Memory (LSTM)

² Stochastic Gradient Descent (SGD)

عمیق‌تری از فرایندهای هیدرولوژیکی ارائه شود. یکی از مهم‌ترین دستاوردهای این تحقیق، توضیح ۶۹ درصد و ۷۸ درصد واریانس برای غلظت رسوب معلق و بار بستر بود که اهمیت مهندسی ویژگی را در مدل‌سازی نشان می‌دهد. ابزارهای به‌کاررفته شامل XGBoost و ویژگی‌هایی مانند جریان نرمال‌شده و شیب هیدروگراف بودند، و همچنین از SHAP⁵ برای تفسیر نتایج استفاده شد. هرچند مقایسه مستقیم با مدل‌های دیگر انجام نشد، اما نتایج نشان داد که XGBoost می‌تواند در دقت پیش‌بینی از مدل‌های سنتی پیشی بگیرد.

در پژوهشی Mohammadi et al., (2025) به پیش‌بینی بار رسوب معلق (SSL) روزانه در حوضه طالقان با استفاده از شش مدل یادگیری ماشین (ML) شامل Cubist، xgbTree و qmn پرداختند. داده‌های ورودی شامل دبی، رسوب روزانه و بارش بین سال‌های ۲۰۰۰ تا ۲۰۱۸ مورد استفاده قرار گرفتند و عملکرد مدل‌ها با معیارهای مختلف ارزیابی شد. نتایج نشان داد که مدل‌های یادگیری ماشین، به‌ویژه xgbTree، بهترین دقت را در پیش‌بینی SSL ارائه دادند و می‌توانند به‌عنوان ابزاری موثر برای تحلیل دینامیک رسوب در حوزه‌های آبخیز به‌کار روند.

نوآوری این تحقیق در استفاده از مدل یادگیری جمعی XGBoost و ترکیب آن با تکنیک‌های خوشه‌بندی و بهینه‌سازی داده‌ها برای بهبود دقت پیش‌بینی‌ها است. همچنین، این مطالعه به بررسی شرایط خاص هیدرولوژیکی و چولگی داده‌ها در این منطقه می‌پردازد. نتایج این تحقیق می‌تواند به‌عنوان یک کار الگویی در شبیه‌سازی رسوب معلق در دیگر ایستگاه‌های هیدرومتری کشور مورد استفاده محققین و کارشناسان این حوزه قرار گیرد. به‌طور کلی، این پژوهش نه‌تنها به درک بهتر فرایندهای رسوگذاری کمک می‌کند، بلکه می‌تواند زمینه‌ساز تحقیقات آتی در جهت بهبود و توسعه مدل‌های پیش‌بینی سازگارتر با شرایط متغیر محیطی باشد.

از یادگیری عمیق بررسی شد تا الگوهای زمانی در رودخانه‌های پرسوب شبیه‌سازی شود. مدل‌های مورد استفاده شامل LSTM و Random Forest بودند. مقایسه‌ها نشان داد که LSTM در پیش‌بینی‌های زمانی عملکرد بهتری نسبت به Random Forest دارد. در پژوهشی که توسط Khosravi et al., (2023) در نهر Esopus در نیویورک، ایالات متحده انجام شد، مدل‌سازی غلظت رسوب معلق روزانه با استفاده از تکنیک‌های یادگیری عمیق بررسی شد تا برآورد دقیق‌تری در رودخانه‌های با جریان متغیر ارائه شود. یکی از مهم‌ترین دستاوردهای این تحقیق، افزایش دقت پیش‌بینی از طریق انتخاب ویژگی‌های مناسب، نظیر جریان و رسوب بالادست/پایین‌دست بود که با داده‌های USGS ارزیابی شد.

ابزارهای به‌کاررفته شامل شبکه‌های عصبی LSTM، تحلیل داده‌های اکتشافی و مهندسی ویژگی بودند. اگرچه این مدل‌ها به‌صورت مستقیم با دیگر مدل‌ها مقایسه نشدند، اما با تمرکز بر LSTM، نشان دادند که یادگیری عمیق می‌تواند الگوهای زمانی پیچیده رسوب را به‌خوبی شبیه‌سازی کند و برای کاربردهای هیدرولوژیکی بسیار مفید باشد. در پژوهشی که توسط Khosravi et al., (2022) در حوزه آبخیز تالار در ایران انجام شد، مدل‌های مختلف برای پیش‌بینی بار رسوب معلق مقایسه شدند تا بهترین رویکرد داده‌محور برای رودخانه‌های پرسوب شناسایی شود. یکی از مهم‌ترین دستاوردهای این تحقیق، دستیابی به دقت بالا در اعتبارسنجی با ضریب ناش¹ در سطح بسیار خوب بود که اهمیت بهینه‌سازی مدل‌ها را به‌خوبی نمایان ساخت. مدل‌های مورد بررسی شامل منحنی سنج رسوب²، جنگل تصادفی³ و مدل SWAT⁴ بودند. مقایسه‌ها نشان داد که مدل هیبریدی dagging-RF (DA-RF) از دقت بالاتری نسبت به SRC و SWAT برخوردار است.

در پژوهشی که توسط Lund et al., (2022) در مینه‌سوتا انجام شد، پیش‌بینی حمل رسوب رودخانه‌ای با استفاده از یادگیری ماشین بهبود یافت تا بینش

⁴ Soil and water assessment tool (SWAT)

⁵ Shapley additive explanation (SHAP) plots

¹ Nash-Sutcliffe Efficiency (NSE)

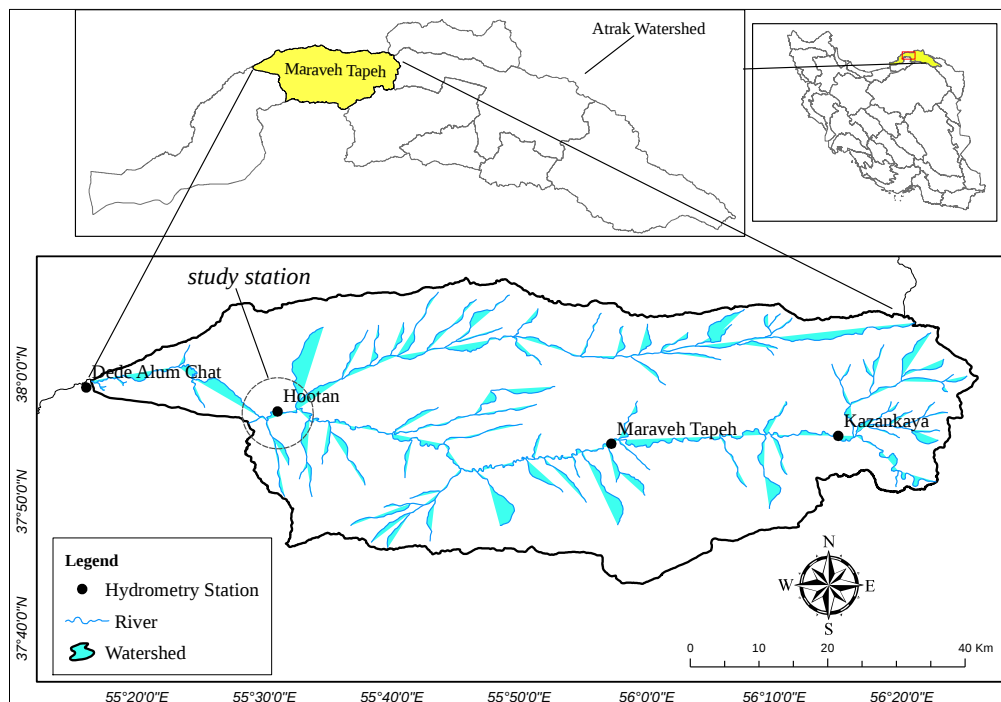
² Sediment rating curve (SRC)

³ Random forest (RF)

مواد و روش‌ها

مطالعاتی مراوه تپه و در محل ایستگاه هیدرومتری هوتن بر روی رودخانه اترک انجام شده است (شکل ۱).

منطقه مورد مطالعه و داده‌های مورد استفاده: پژوهش حاضر در استان گلستان، حوزه آبخیز اترک، واحد



شکل ۱- موقعیت حوزه آبخیز مورد مطالعه و ایستگاه هیدرومتری هوتن

Fig. 1. Location of the studied watershed and the Hootan hydrometric station

و دارای اشل، پل تلفریک و لیمنوگراف است. داده‌های آب‌سنجی مورد استفاده در این پژوهش، شامل دبی لحظه‌ای، دبی جریان روزانه و دبی رسوب معلق روزانه ایستگاه مربوطه بوده که از مرکز آمار تحقیقات منابع آب وزارت نیرو (تماب) و شرکت آب منطقه‌ای استان گلستان اخذ شده است.

دوره آماری طرح بین سال‌های ۱۳۶۳ تا ۱۴۰۰ است. تعداد رکوردهای اطلاعاتی مورد استفاده پس از حذف آمار خالی و پرت مجموعاً ۴۰۶ نمونه بوده است. علاوه بر داده‌های دبی جریان و رسوب معلق، به‌منظور بررسی تأثیر به‌کارگیری داده‌های بارندگی در میزان کارایی مدل‌های شبیه‌ساز رسوب معلق رودخانه اترک، از آمار متوسط بارش روزانه سه ایستگاه اندازه‌گیری بارش واقع در منطقه مورد مطالعه به نام‌های هوتن، مراوه تپه و قازانقایه استفاده شد که خلاصه آماری این داده‌ها نیز در جدول ۱ ارائه شده است.

از نظر اقلیمی حوزه اترک، دارای اقلیم خشک و نیمه‌خشک است و میزان متوسط بارندگی سالانه آن ۳۴۰ میلی‌متر و متوسط درجه گرم‌ترین ماه سال، ۲۵ درجه سانتیگراد است (Etka, 2024). رودخانه اترک به‌دلایل مختلف، از جمله قرار گرفتن روی سازندهای زمین‌شناسی فرسایش‌پذیر و فعالیت‌های انسانی مانند تغییر کاربری اراضی و غیره، دارای تغییرات مورفولوژیکی زیاد است. در بازدید میدانی مشخص شد که فرسایش قهقراپی به سمت بالادست آبراه‌ها به‌صورت فرسایش خندقی و شیاری تا سطح زمین متصل شده و به‌صورت پسرونده‌ای به سمت زمین‌های کشاورزی در حال حرکت مخرب است (شکل‌های ۲ و ۳) (Etka, 2024).

ایستگاه هیدرومتری هوتن با کد ۰۷۳-۱۱، در موقعیت $37^{\circ}57'07''$ طول شرقی و $55^{\circ}31'01''$ عرض شمالی در سال ۱۳۵۶ بر روی رودخانه اترک تأسیس شده است. این ایستگاه به لحاظ کیفیت تجهیزات، جزء ایستگاه‌های درجه ۲ محسوب می‌شود



شکل ۲- فرسایش قهقراپی و توسعه خندق‌ها به سمت بالادست در زمین‌های بایر (شمال غرب ایستگاه هیدرومتری هوتن)

Fig. 2. Regressive erosion and the development of gullies upstream in barren lands (Northwest of the Hootan hydrometric station)



شکل ۳- وضعیت بستر آبراهه متصل شده به رودخانه اترک (شمال غرب ایستگاه هیدرومتری هوتن)

($37^{\circ}59'30.7''N$ $55^{\circ}24'41.8''E$)

Fig. 3. Condition of the streambed connected to the Atrak River (Northwest of the Hootan Hydrometric Station) ($37^{\circ}59'30.7''N$ $55^{\circ}24'41.8''E$)

جدول ۱- ویژگی‌های آماری داده‌های مورد استفاده در شبیه‌سازی رسوب معلق ایستگاه آب‌سنجی هوتن (۱۳۶۳ تا ۱۴۰۰)

Table 1. Statistical characteristics of the data used in the simulation of suspended sediment at the Hootan hydrometric station (1984 to 2021)

Data Type	Statistical Parameters						
	Mean	Med.	Std.	Max.	Min.	Skew.	C.V%
Suspended sediment concentration (mg l^{-1})	17319.21	1164.66	41446.42	350530.7	6	4.15	239.309
Suspended sediment load (tonday^{-1})	41883.70	414.52	184357.437	1713188.73	0.016	678	440.16
Daily Flow discharge (m^3s^{-1})	9.52	4.13	18.55	173	0	4.67	184.83
Instantaneous flow discharge (m^3s^{-1})	8.74	4.08	16.98	141.20	0.004	4.79	194.18
Daily precipitation (mm)	0.86	0	2.23	14.33	0	3.46	257.69

حداکثر آن بسیار زیاد است. وجود انحراف معیار زیاد داده‌های رسوب معلق در مقایسه با میانگین آن، حکایت از غیرنرمال بودن و چولگی (به راست) زیاد

همانطور که از اطلاعات آماری جدول ۱ استنباط می‌شود، دبی رسوب دارای چولگی و ضریب تغییرات بسیار زیاد بوده و همچنین تغییرات بین حداقل و

هر نوع تابع غیرخطی هستند (Hornik et al., 1989). در پژوهش حاضر، از یک شبکه عصبی پرسپترون سه لایه (یک لایه ورودی، یک لایه پنهان و یک لایه خروجی) استفاده شده و تعداد بهینه نورون‌ها با سعی و خطا تعیین شد. به منظور آموزش شبکه عصبی از الگوریتم یادگیری Adam به دلیل کارایی و همگرایی سریع‌تر در آموزش شبکه و اثبات کارآمدی آن در تخمین رسوب رودخانه‌ای استفاده شده است (Sahoo et al., 2022; Kaveh et al., 2019). در آموزش مدل‌های شبکه عصبی، توابع فعال‌سازی مختلفی نظیر Sigmoid و tanh در نورون‌های شبکه مورد استفاده قرار گرفت.

یادگیری عمیق با شبکه عصبی متصل متراکم^۵: یادگیری عمیق، زیرشاخه‌ای از یادگیری ماشین است که از شبکه‌های عصبی مصنوعی^۶ برای نگاشت ویژگی‌ها به خروجی‌ها استفاده می‌کند (Ulke et al., 2016). شبکه‌های عصبی DNN، RNN^۷ و LSTM^۸ از انواع این شبکه‌ها هستند. در DNN، تمام گره‌های لایه قبلی به گره‌های لایه بعدی متصل‌اند و نسخه اولیه آنها پرسپترون چندلایه نامیده می‌شود. این شبکه‌ها با یک لایه ورودی، لایه‌های پنهان و یک لایه خروجی، رابطه‌های خطی و غیرخطی بین ورودی‌ها و خروجی‌ها را یاد می‌گیرند (Sit et al., 2020). در این پژوهش، از DNN برای شبیه‌سازی رسوب معلق رودخانه اترک استفاده شده است. انتخاب الگوریتم بهینه‌سازی مناسب برای آموزش مدل یادگیری عمیق از اهمیت بالایی برخوردار است و تأثیر زیادی بر کیفیت نتایج دارد. الگوریتم آدام^۷، نسخه تعمیم‌یافته‌ای از «گرادیان نزولی تصادفی»^۸ است که به‌طور گسترده‌ای در کاربردهای یادگیری عمیق، به‌ویژه در بینایی کامپیوتر و پردازش زبان طبیعی، به کار می‌رود.

در پژوهش حاضر به‌منظور آموزش شبکه عصبی از این الگوریتم به دلیل کارایی و همگرایی سریع‌تر در آموزش شبکه و اثبات کارآمدی آن در تخمین رسوب رودخانه‌ای استفاده شده است (Sahoo et al., 2022).

این داده‌ها دارد. این موضوع، به همراه سایر آماره‌های محاسبه شده، بیانگر پیچیدگی مدل‌سازی و برآورد رسوب معلق رودخانه است.

شبکه عصبی مصنوعی پرسپترون چند لایه روبه جلو^۱: شبکه عصبی مصنوعی پرسپترون چند لایه، از عناصر عملیاتی ساده‌ای به نام نورون ساخته شده که به صورت موازی و در کنار هم، عمل می‌کنند. شبکه‌های عصبی چند لایه، از یک لایه ورودی (در ارتباط با داده‌های ورودی)، یک یا چند لایه پنهان (جهت سازمان‌دهی نورون‌ها) و یک لایه خروجی (در رابطه با داده خروجی) تشکیل می‌شوند. عملکرد شبکه عصبی، از طریق نحوه اتصال بین اجزاء، با تنظیم مقادیر هر اتصال که به نام وزن اتصال بیان می‌شود، تعیین می‌شود. یکی از انواع شبکه‌های عصبی پرکاربرد در هیدرولوژی و منابع آب، شبکه‌های عصبی پرسپترون چند لایه روبه جلو با الگوی آموزش پس انتشار^۲ خطا است (Ulke et al., 2023).

در این نوع از شبکه‌های عصبی، جهت جریان داده، از لایه ورودی به سمت لایه پنهان و از لایه پنهان به سمت لایه خروجی بوده و از این نظر به آنها شبکه‌های عصبی روبه جلو یا پیش‌خور گفته می‌شود. به‌منظور آموزش شبکه عصبی، مقدار خطا در جهت بیشترین شیب تابع خطا محاسبه شده و این مقدار، به لایه‌های قبل (لایه یا لایه‌های پنهان) فرستاده شده تا با تنظیم مجدد مقادیر وزن نورون‌ها، مقدار خطا را کاهش دهند (قانون دلتا) (Tayfur et al., 2012) (رابطه ۱).

$$w_{ij}^{new} = w_{ij}^{old} - \eta \frac{\partial E}{\partial w_{ij}} \quad (1)$$

که در آن، w_{ij}^{old} و w_{ij}^{new} به ترتیب وزن بین نورون‌های i و j قبل و بعد از یک تکرار معین، η نرخ یادگیری و E تابع خطا است.

آموزش شبکه و کاهش خطا تا ایجاد همگرایی در شبکه ادامه می‌یابد. شبکه‌های عصبی می‌توانند دارای چندین لایه پنهان باشند با این وجود، تحقیقات انجام شده نشان می‌دهد که شبکه‌های عصبی پیش‌خور، با دارا بودن یک لایه پنهان، قادر به تقریب زدن

⁵ Recurrent Neural Networks

⁶ Long Short Term Memory Networks

⁷ Adam Optimization Algorithm (Adam)

⁸ Stochastic Gradient Descent (SGD)

¹ Feed-forward Multi-layer Perceptron (FFMLP)

² Back Propagation

³ Densely Connected Neural Network (DNN)

⁴ Artificial Neural Network (ANN)

مدل از مجموع پیش‌بینی هر یادگیرنده به‌دست می‌آید (Chen and Guestrin, 2023).

در این پژوهش، برای تنظیم ابرپارامترهای مدل XGBoost، از روش Random Search استفاده شده است. این روش این امکان را می‌دهد که در فضای بزرگ ابرپارامترها به‌طور کارآمد به جستجوی مقادیر بهینه پرداخت، به‌ویژه برای پارامترهایی مانند نرخ یادگیری و عمق درخت که تأثیر زیادی بر عملکرد مدل دارند. تنظیم ابرپارامترها به‌طور خاص شامل انتخاب مقادیر بهینه برای نرخ یادگیری (learning rate) است که بر سرعت یادگیری مدل تأثیر می‌گذارد و عمق درخت (max depth) که بر پیچیدگی مدل و توانایی آن در یادگیری روابط غیرخطی تأثیر دارد. این تنظیمات به بهبود دقت مدل و کاهش خطر تطبیق بیش از حد کمک می‌کند و به‌طور کلی کیفیت پیش‌بینی‌ها را افزایش می‌دهد.

شبکه عصبی نگاشت خود سازمان ده: شبکه عصبی نگاشت خود سازمان‌ده یا SOM^۶، یک شبکه عصبی مصنوعی غیرنظارتی بوده و الگوریتم آموزش آن، به‌صورت رقابتی و بدون ناظر انجام می‌شود. ساختار SOM، از یک لایه ورودی و یک لایه خروجی (لایه کوهنن) تشکیل می‌شود (Kohonen, 2002). در این ساختار، نورون‌های لایه ورودی، محل ارتباط داده‌های ورودی با شبکه بوده و به ازاء هر متغیر ورودی، یک نورون در این لایه وجود دارد (به‌عنوان مثال دبی، بارش و غیره). لایه خروجی، ماتریس یا شبکه‌ای (عموماً یک شبکه دو بعدی) از نورون‌ها را تشکیل داده به‌نحوی که هر نورون این شبکه، به کلیه نورون‌های لایه ورودی متصل بوده ولی به نورون‌های دیگر این لایه متصل نیست. فرایند آموزش در شبکه SOM از سه مرحله رقابت^۷، همکاری^۸ و تطبیق^۹ تشکیل می‌شود. در ابتدا وزن نورون‌های لایه خروجی به‌صورت تصادفی تعریف شده و در طی فرایند آموزش و در طول زمان، به مقادیر متغیرهای بردارهای ورودی بیشتر شبیه می‌شود. پس از آنکه شبیه‌ترین نورون به

(Kaveh et al., 2019). علاوه بر آن، تعداد لایه پنهان و تعداد بهینه نورون‌ها با سعی و خطا تعیین شد. از توابع فعال‌سازی ReLU، Leaky ReLU، Sigmoid و Tanh در نورون‌ها و از توابع زیان^۱ MSE، MAE و Hubber برای بهینه‌سازی ماتریس ضرایب شبکه‌های عصبی استفاده شد. همچنین، به‌منظور جلوگیری از بیش‌برازش^۲ مدل‌ها در مدت زمان آموزش، از الگوریتم توقف زودرس^۳ استفاده شد.

یادگیری جمعی^۴: الگوریتم XGBoost^۵، توسط Chen et al., (2016) پیشنهاد شده است. این الگوریتم بر پایه درخت‌های تصمیم بنا شده و از تکنیک تقویت گرادیان برای دستیابی به دقت و کارایی بی‌نظیر بهره می‌برد. الگوریتم XGBoost یک الگوریتم یادگیری با ناظر است که به‌خاطر قابلیت‌هایش در تنظیم خودکار، سرعت بالا در آموزش و قدرت پیش‌بینی‌اش، به‌عنوان یکی از قدرتمندترین و مؤثرترین ابزارها در حوزه یادگیری ماشین شناخته می‌شود و در طیف گسترده‌ای از کاربردهای عملی، از تجزیه و تحلیل داده‌های مالی گرفته تا پیش‌بینی‌های پزشکی، استفاده می‌شود. این الگوریتم مبتنی بر ایده "تقویت" است که تمام پیش‌بینی‌های یادگیرندگان "ضعیف" را برای ایجاد یک یادگیرنده "قوی" از طریق استراتژی‌های آموزش ترکیب می‌کند.

هدف XGBoost جلوگیری از تطبیق بیش از حد و همچنین بهینه‌سازی منابع محاسباتی است. این امر با ساده‌سازی توابع هدف به‌دست می‌آید که امکان ترکیب شرایط پیش‌بینی و منظم‌سازی را فراهم می‌کند، اما سرعت محاسباتی بهینه را حفظ می‌کند. همچنین، محاسبات موازی به‌طور خودکار در مرحله آموزش برای توابع در XGBoost اجرا می‌شود. یادگیرنده اول، ابتدا به کل فضای داده‌های ورودی برازش داده می‌شود و سپس مدل دوم برای رفع اشکال‌های یک یادگیرنده ضعیف برازش داده می‌شود. این فرایند برازش برای چند بار تکرار می‌شود تا زمانی که معیار توقف برآورده شود. پیش‌بینی نهایی

^۶ Self-Organizing Map (SOM)

^۷ Competitive Phase

^۸ Co-operative Phase

^۹ Adaptive Phase

^۱ Loss functions

^۲ Overfitting

^۳ Early stopping method

^۴ Ensemble learning

^۵ Extreme Gradient Boosting

لازم است داده‌های آزمون، مشابه داده‌های آموزش و ارزیابی متقاطع^۳ (که برای جلوگیری از بیش برآش مدل استفاده می‌شود) بوده و از توزیع فراوانی یکسانی با آنها برخوردار باشند. بدین منظور در پژوهش حاضر از شبکه عصبی نگاشت خود سازمان‌ده^۴، برای خوشه‌بندی داده‌ها و از روش تخصیص برابر^۵، به منظور نمونه‌گیری از خوشه‌ها استفاده شده است. تعیین تعداد بهینه خوشه‌ها، توسط شاخص ارزیابی کیفیت خوشه‌بندی دیویس-بولدین انجام گرفته است. به منظور تحلیل نتایج آماری حاصل از خوشه‌بندی داده‌ها در سه مجموعه داده آموزش، ارزیابی متقاطع و آزمون، پارامترهای آماری (میانگین، انحراف معیار، چولگی) آنها با یکدیگر مقایسه شد.

نرمال‌سازی خطی داده‌ها: نرمال‌سازی داده‌ها، به منظور بی‌بعد نمودن آنها در محاسبات فاصله (در عملیات خوشه‌بندی)، امکان استفاده از متغیرهای چندگانه با واحدهای اندازه‌گیری متفاوت در یک مدل و یا به منظور جلوگیری از کوچک‌شدن بیش از حد وزن‌های تخصیص یافته به نورون‌ها (در مدل‌های شبکه عصبی) استفاده می‌شود. در این پژوهش، با توجه به استفاده از توابع محرک سیگموئید و تانژانت هایپربولیک از رابطه‌های (۴) و (۵) به ترتیب، برای نرمال نمودن داده‌ها در بازه‌های $[0/1-0/9]$ و $[-0/9-0/9]$ استفاده شده است.

$$z = 0.1 + (0.8 * \frac{X_i - X_{imin}}{X_{imax} - X_{imin}}) \quad (4)$$

$$z = \left(\frac{1.8(X_i - X_{imin})}{X_{imax} - X_{imin}} \right) - 0.9 \quad (5)$$

که در آن، Z ، متغیر نرمال شده، X_i ، متغیر اولیه، X_{imax} و X_{imin} به ترتیب مقادیر کمینه و بیشینه متغیر X_i هستند.

روش نمونه‌گیری از خوشه‌ها: به منظور تهیه سه مجموعه حتی‌الامکان مشابیه و همگن (مجموعه‌های آموزش، ارزیابی متقاطع و آزمون)، بایستی از خوشه‌های تولیدی به شکل مناسبی نمونه‌گیری به عمل آمده که در پژوهش حاضر، از روش تخصیص متناسب (رابطه ۶)، به نسبت ۱۵، ۷۰ و ۱۵ درصد

بردار ورودی^۱ مشخص شد، وزن‌های آن و وزن‌های دیگر نورون‌های همسایه آن، برحسب مقدار فاصله‌ای که از BMU دارند (فاز همکاری) بر طبق رابطه (۲) به هنگام می‌شوند (فاز تطبیق).

$$wji(t+1) = wji(t) + \theta(t) * \eta(t) * [xi(t) - wji(t)] \quad (2)$$

که در آن، t ، نماینده زمان، $\theta(t)$ ، تابعی است که فاصله نورون‌های همسایه BMU را به نسبتی از همسایگی تبدیل می‌کند و $\eta(t)$ ، نرخ یادگیری است. در پژوهش حاضر همانطور که پیشتر بیان شد، به منظور تعیین تعداد خوشه‌های بهینه، از شاخص ارزیابی دیویس-بولدین^۲ استفاده شده است (Kaufman and Rousseeuw, 2009).

مدل رگرسیونی منحنی سنج رسوب: مدل رگرسیون منحنی سنج رسوب، یک رابطه توانی بین دبی رسوب (یا غلظت رسوب) و دبی جریان است. Jones et al., (1981) رابطه (۳) را برای این مدل پیشنهاد نمودند.

$$Q_s = aQ^b \quad \text{or} \quad \log Q_s = b \log Q + \log a \quad (3)$$

که در آنها، Q دبی جریان (مترمکعب بر ثانیه)، Q_s دبی رسوب معلق (تن در روز یا کیلوگرم بر ثانیه)، C غلظت رسوب معلق (میلی گرم در لیتر یا کیلوگرم در مترمکعب) و a و b ضرائب ثابت رابطه رگرسیون هستند.

در روش حد وسط دسته‌ها که آن را Jansson (1996) پیشنهاد کرده است، دبی جریان، به تعدادی دسته یا طبقه تقسیم شده (ارتفاع دسته‌ها یا طبقات یکسان است) و برای دبی متوسط هر دسته، متوسط دبی رسوب اندازه‌گیری شده همان دسته تعیین می‌شود و پس از برآش یک مدل خطی (همانند توضیحات بیان شده در قسمت قبل)، ضرائب منحنی سنج رسوب با استفاده روش حداقل مربعات خطا به دست می‌آید.

تهیه داده‌های همگن جهت آموزش و آزمون مدل‌ها: به منظور افزایش قدرت تعمیم‌دهی مدل‌های داده مبنا، لازم است داده‌های مورد استفاده در آموزش مدل، معرف و نماینده داده‌های کل دوره آماری باشند. همچنین به منظور آزمون و ارزیابی صحیح نتایج مدل،

⁵ Proportional Allocation

¹ Best Matching Unit (BMU)

² Davies-Bouldin Index

³ Cross validation

⁴ Self-Organizing Map (SOM)

تپه و قازانقایه) از روز جاری تا پنج روز قبل آن است. به‌منظور انتخاب بهینه متغیرهای ورودی (پیش‌بینی کننده) در مدل‌های هوشمند، از امکانات موجود در مدل جنگل تصادفی^۱ استفاده شد.

صحت‌سنجی و ارزیابی مدل‌ها: به‌منظور ارزیابی و مقایسه نتایج گرفته شده از مدل‌های داده مبنا و مقایسه آنها با داده‌های رسوب مشاهداتی (داده‌های مجموعه آزمون)، اندازه‌گیری مقدار خطا و ترسیمات گرافیکی انجام شد. در این رابطه از شاخص‌های ریشه میانگین مربعات خطا^۲، میانگین قدر مطلق خطا (MAE)، ضریب کارائی ناش-ساتکلیف^۳، ضریب تبیین^۴ و درصد اریبی یا بایاس (PBIAS) بین داده‌های مشاهده‌ای با داده‌های محاسباتی (به‌ترتیب رابطه‌های ۹ تا ۱۳) استفاده شد.

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (S_M - S_O)^2} \quad (9)$$

$$MAE = \frac{\sum_{i=1}^n |S_O - S_M|}{n} \quad (10)$$

$$NS = 1 - \frac{\sum_{i=1}^n (S_M - S_O)^2}{\sum_{i=1}^n (S_O - \bar{S}_O)^2} \quad (11)$$

$$R^2 = \left[\frac{\sum_{i=1}^n (S_O - \bar{S}_O)(S_M - \bar{S}_M)}{\sqrt{\sum_{i=1}^n (S_O - \bar{S}_O)^2 \sum_{i=1}^n (S_M - \bar{S}_M)^2}} \right]^2 \quad (12)$$

$$PBIAS = \frac{\sum_{i=1}^n (S_O - S_M)}{\sum_{i=1}^n (S_O)} * 100 \quad (13)$$

در رابطه‌های فوق، S_O و S_M به‌ترتیب، غلظت رسوب معلق مشاهده‌ای و برآورد شده بر حسب میلی‌گرم در لیتر، $\bar{S}_O$ و $\bar{S}_M$ ، به‌ترتیب میانگین غلظت رسوب مشاهده‌ای و محاسبه شده و n ، تعداد داده‌ها است.

پیش از ارائه نتایج، به‌منظور درک بهتر از روش‌های بیان شده در فرایند مدلسازی رسوب معلق و چگونگی ارتباط بین بخش‌های مختلف پژوهش، ساختار کلی مراحل مدلسازی، در شکل ۴ خلاصه و نمایش داده شده است.

به‌ترتیب برای سه مجموعه داده یاد شده استفاده شده است. در روش تخصیص متناسب، تعداد نمونه‌ها متناسب با اندازه خوشه تغییر کرده بدین نحو که با افزایش اندازه هر خوشه، تعداد نمونه‌گیری از آنها افزایش و کاهش یافته است (رابطه ۶).

$$nh = n \frac{N_h}{\sum_{j=1}^H N_j} \quad (6)$$

که در آن، nh ، تعداد نمونه گرفته شده از خوشه h ، n ، تعداد داده مورد نیاز، N_h ، تعداد داده‌ها در خوشه h و N_j ، تعداد داده در سایر خوشه‌ها است.

در پژوهش حاضر، به‌منظور پیاده‌سازی مدل‌های شبکه عصبی، خوشه‌بندی داده‌ها و همچنین انجام تحلیل‌های آماری، از نرم‌افزارهای پایتون، C# و سامانه مکانی هوشمند شبیه‌ساز رسوب معلق که توسط نگارنده تهیه و تدوین شده و دارای گواهی ثبت اختراع به شماره ۹۹۹۲۵ است، استفاده شده است.

الگوی داده‌ها به منظور شبیه‌سازی رسوب معلق رودخانه اترک: به‌منظور شبیه‌سازی رسوب معلق رودخانه اترک (در محل ایستگاه هیدرومتری هوتن) از مدل‌های رگرسیونی (منحنی سنج‌های رسوب تک خطی و حد وسط دسته‌ها) و مدل‌های هوشمند (شبکه‌های عصبی، یادگیری عمیق و یادگیری جمعی) استفاده خواهد شد. لذا، الگوی کلی متغیرهای ورودی و خروجی به‌صورت رابطه‌های (۷) و (۸) به‌ترتیب برای مدل‌های هوشمند و رگرسیونی در زیر هستند.

$$SSL = f(Inst.D, D_0, D_1, D_2, D_3, D_4, D_5, P0, P1, P2, P3, P4, P5) \quad (7)$$

$$SSL = f(Inst.D) \quad (8)$$

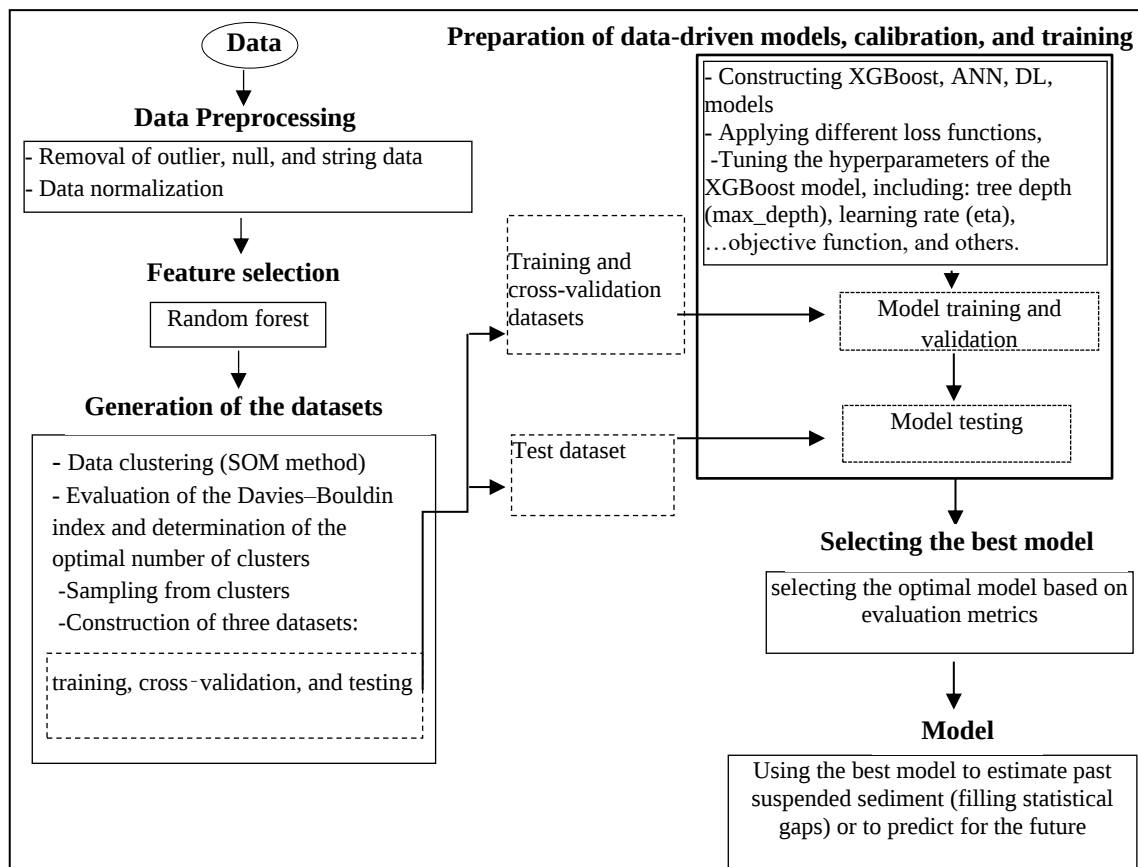
که در آن، SSL ، بار رسوب معلق روزانه (تن در روز)، $Inst.D$ ، دبی لحظه‌ای جریان (مترمکعب در ثانیه)، D_0 تا D_5 ، دبی‌های روزانه روز جاری تا ۵ روز قبل و $P0$ تا $P5$ ، میانگین بارش ۳ ایستگاه بارندگی (هوتن، مراوه

³ Nash-Sutcliffe model efficiency coefficient (NSE)

⁴ Coefficient of Determination (R^2)

¹ Random forest

² Root Mean Square Error (RMSE)



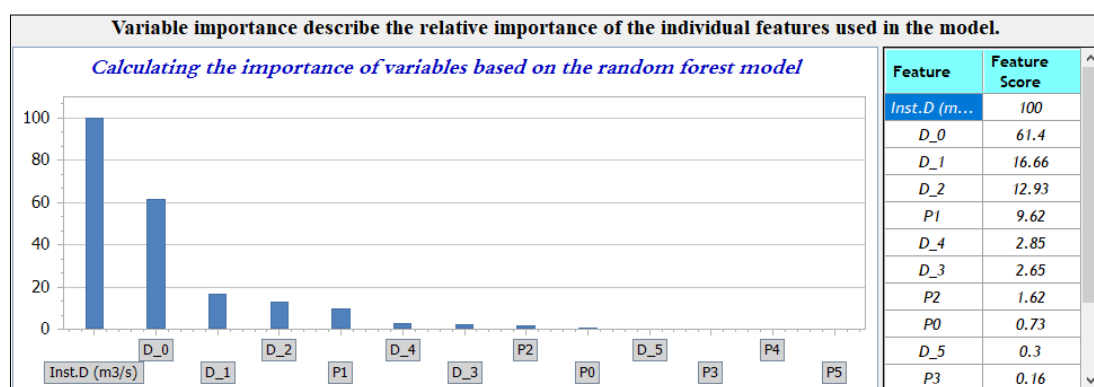
شکل ۴ - روندنمای مراحل مختلف پژوهش

Fig. 4. Flowchart of the various stages of the research

نتایج و بحث

مترمکعب در ثانیه)، D_5 تا D_0، دبی‌های روزانه روز جاری تا پنج روز قبل و P_5 تا P_0، میانگین بارش روز جاری تا پنج روز قبل) با استفاده از الگوریتم جنگل تصادفی در شکل ۵ نشان داده شده است.

نتایج انتخاب ویژگی: نتایج گرفته شده در رابطه با انتخاب ویژگی (متغیرهای پیش‌بینی کننده) از بین ویژگی‌های مورد مطالعه (Inst.D، دبی لحظه‌ای جریان



شکل ۵- نتایج انتخاب ویژگی با استفاده از الگوریتم جنگل تصادفی

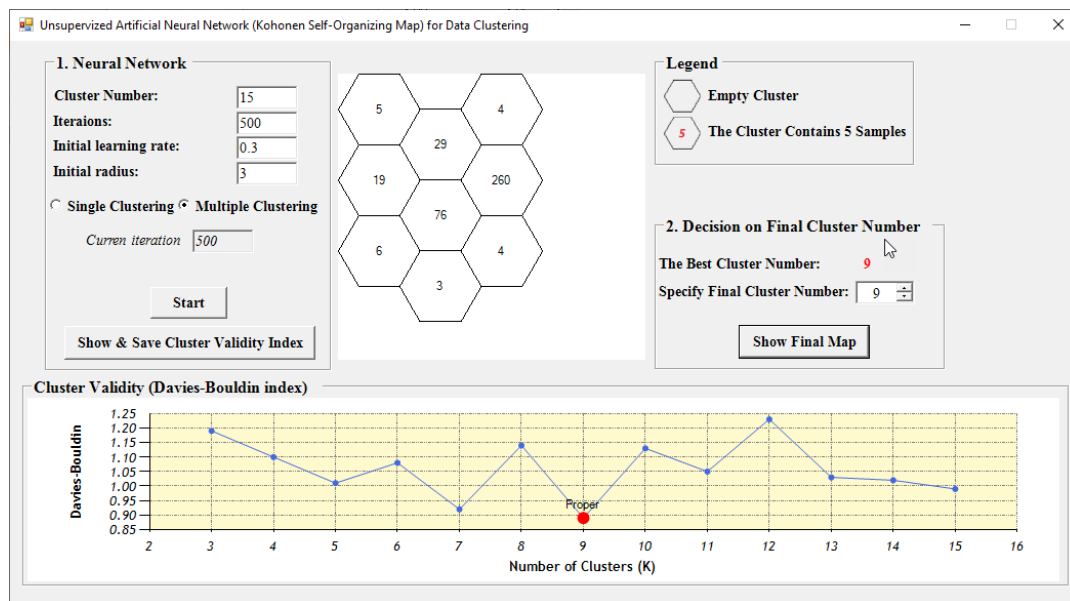
Fig. 5. Feature selection results using the random forest algorithm

یک روز قبل (P_1) در شبیه‌سازی رسوب معلق رودخانه اترک استفاده شد.

با توجه به امتیازات محاسبه شده برای ویژگی‌ها (مقادیر Feature Score)، از ویژگی‌های دبی لحظه‌ای (Inst.D)، دبی‌های روزانه D_0 تا D_2 و میانگین بارش

مطالعه ایستگاه هیدرومتری هوتن تعداد نه عدد محاسبه شد (شکل ۶).

نتایج خوشه‌بندی و تحلیل آماری داده‌ها: با استفاده از خوشه‌بندی SOM و محاسبه شاخص دیویس-بولدین، تعداد بهینه خوشه‌ها برای مجموعه داده‌های تحت



شکل ۶- تعیین تعداد بهینه خوشه‌ها با استفاده از خوشه‌بندی SOM و شاخص دیویس-بولدین در ایستگاه هیدرومتری هوتن
Fig. 6. Determination of the optimal number of clusters using som clustering and the Davies-Bouldin index at the Hootan Hydrometric Station

سعی شده تا با کمک خوشه‌بندی و نمونه‌گیری از آنها، ویژگی‌های آماری متغیرهای ورودی و خروجی تا حد امکان شبیه به هم باشند. همچنین سعی بر این بوده تا مقادیر حدی ویژگی‌ها، در مجموعه‌های آموزش و پس از آن در مجموعه اعتبارسنجی متقابل قرار گیرد. در این رابطه همانطور که در جدول ۲ آمده است، بیشترین مقادیر دبی لحظه‌ای و دبی رسوب معلق به ترتیب ۱۴۱/۲۰ مترمکعب در ثانیه و ۱۷۱۳۱۸۸/۷۳ تن در روز متعلق به مجموعه داده‌های آموزش هستند.

پس از تعیین تعداد بهینه خوشه‌ها، به نسبت ۷۰، ۱۵، ۱۵ درصد (به ترتیب برای مجموعه‌های داده آموزش، اعتبارسنجی متقاطع و آزمون)، از خوشه‌ها نمونه‌گیری به عمل آمده که نتایج خلاصه پارامترهای آماری (معیارهای مرکزیت و پراکندگی داده‌ها) ویژگی‌های منتخب در سه مجموعه داده‌های آموزش، اعتبارسنجی متقاطع و آزمون برای ایستگاه هیدرومتری هوتن در جدول ۲ آمده است. همانطور که مقادیر پارامترهای محاسبه شده در جدول ۲ نشان می‌دهد،

جدول ۲- پارامترهای آماری مجموعه داده‌های آموزش، اعتبارسنجی متقاطع و آزمون ایستگاه هیدرومتری هوتن

Table 2. Statistical parameters of the training, cross-validation, and testing datasets at the Hootan hydrometric station

Model Variables and Dataset	Statistical Parameters						
	Max.	Min.	Mean	Median	SD	Kurt.	CV%
Instantaneous Flow Discharge (Inst.D) (m ³ /s)							
Training Set	141.2	0.01	8.48	4.11	16.69	5.07	196.86
Cross-Validation Set	124.61	0.01	10.47	4.08	20.71	4.04	197.88
Test Set	86.22	0	8.29	4	13.96	3.65	168.46

Daily Flow Discharge (D_0) (m ³ /s)							
Training Set	128	0	9	4.35	17.08	4.41	189.73
Cross-Validation Set	173	0	11.62	4.11	25.7	4.71	221.14
Test Set	86.1	0	9.79	4.02	16.34	3.16	166.85
Discharge with a 1-Day Lag (D_1) (m ³ /s)							
Training Set	171	0	9.17	4	19.58	5.14	213.43
Cross-Validation Set	83.3	0	7.81	4.17	12.36	3.96	158.32
Test Set	99.8	0	10.1	4.16	18.95	3.62	187.68
Discharge with a 2-Day Lag (D_2) (m ³ /s)							
Training Set	120	0	7.63	4.02	13.74	4.96	180.08
Cross-Validation Set	39.5	0	6.91	3.3	8.85	1.7	128.16
Test Set	71.8	0	8.67	4.73	13.69	3.32	157.81
Average Precipitation with a 1-Day Lag (P_1) (mm)							
Training Set	14.33	0	0.83	0	2.18	3.53	263.38
Cross-Validation Set	13	0	0.9	0	2.31	3.34	256.29
Test Set	13.33	0	1.01	0	2.39	3.23	236.9
Suspended Sediment Load (SSL) (ton/day)							
Training Set	1713188.73	0.02	39608.72	416.91	184556.82	6.76	465.95
Cross-Validation Set	1709084.53	0.14	64999.11	408.91	240635.44	5.62	370.21
Test Set	630153.57	0.04	29840.44	391.75	100950.84	4.75	338.3

در جدول‌های ۳ و ۴، ویژگی‌ها و مشخصات مدل‌های یادگیری عمیق، شبکه عصبی چندلایه MLP و XGBoost و مقادیر بهینه آنها آمده است.

جدول ۳- مشخصات مدل‌های یادگیری عمیق و شبکه عصبی چندلایه (MLP)

Table 3. Specifications of deep learning models and Multi-Layer Perceptron (MLP)

Feature	MLP Model	Deep Learning Model
Number of Layers	2 (1 hidden layer + 1 output layer)	7
Hidden Layers	1	5 layers with 15 neurons each
Output Layer	1 neuron	1 neuron
Dropout Rate	-	0.2
Activation Function	Sigmoid	tanh
Loss Function	Mean Squared Error (MSE)	Huber
Optimization Algorithm	Levenberg-Marquardt	Adam
Number of Iterations	150-200	20-40
Training Method	Levenberg-Marquardt	-
Input Data Type	Numeric Data	Numeric Data
Output Type	Numeric	Numeric
Preprocessing Techniques	Normalization, Log Transformation	Normalization, Log Transformation
Patience (Early Stopping)	7 (Cross Validation)	7 (Cross Validation)

جدول ۴- مشخصات مدل XGBoost
Table 4. XGBoost model specifications

Feature	Details
Model Type	XGBoost Regressor
Hyperparameters	- max_depth =5 - learning_rate =0.1 - colsample_bytree =0.6 - reg_alpha =0.5 - min_child_weight =3 - n_estimators =150 - subsample =0.75 - reg_lambda =0.1 - tweedie_variance_power=1.7
Optimization Algorithm	Gradient Boosting (Tree-based model)
Loss Function	Tweedie (with 'mphe' as the evaluation metric)
Cross-Validation	Repeated K-Fold (10 splits, 3 repeats)
Early Stopping	Yes (with early_stopping_rounds parameter=10)
Performance Metrics	- Train/Test MSE - Train/Test MAE - Train/Test R ²

نتایج ارزیابی و صحت‌سنجی مدل‌ها: نتایج ارزیابی و صحت‌سنجی مدل‌های مختلف داده مینا برای مجموعه داده‌های آزمون، در جدول ۵ و نمودار پراکنش داده‌های مشاهداتی آزمون در مقابل داده‌های تخمینی در مدل‌های مختلف در شکل ۷ نشان داده شده است. در جدول ۳ در مدل‌های شبکه عصبی، تعداد بهینه نرون‌های لایه پنهان با آزمون و خطا و کمینه نمودن

خطای شبکه در تکرارهای زیاد تعیین شده است. همچنین به‌منظور بهینه نمودن مدل‌های شبکه عصبی و مدل یادگیری جمعی، از انواع توابع فعال‌سازی و زیان، استفاده شد. لازم به ذکر مجدد است که آموزش، اعتبارسنجی و آزمون کلیه مدل‌ها با داده‌های یکسان (مجموع داده‌های آموزش، ارزیابی متقاطع و آزمون) انجام شده است.

جدول ۵- نتایج ارزیابی و صحت‌سنجی مدل‌های داده مینا در برآورد بار رسوب معلق ایستگاه هیدرومتری هوتن (برای مجموعه داده‌های آزمون)

Table 5. Results of evaluation and validation of data-driven models in estimating suspended sediment load at the Hootan hydrometric station (for the testing dataset)

Model Name	Validation Indices and Model Performance Evaluation				
	MAE (ton/day)	RMSE (ton/day)	NSE	R2	PBIAS (%)
Neural Network Model (MLP)	16015.53	45052.83	0.8	0.88	-28.66
Deep Learning Neural Network Model (DL)	15979.02	40387.12	0.84	0.84	2.04
Ensemble Learning Model (XGBoost)	10652.13	36219	0.87	0.87	-3.52
Regression Model for Sediment Rating Curve (Logged Mean Load within Discharge Class)	20909.23	45632.08	0.8	0.81	-18.93
Regression Model for Sediment Rating Curve (Ordinary linear regression)	24048.48	84509.38	0.3	0.8	76.97

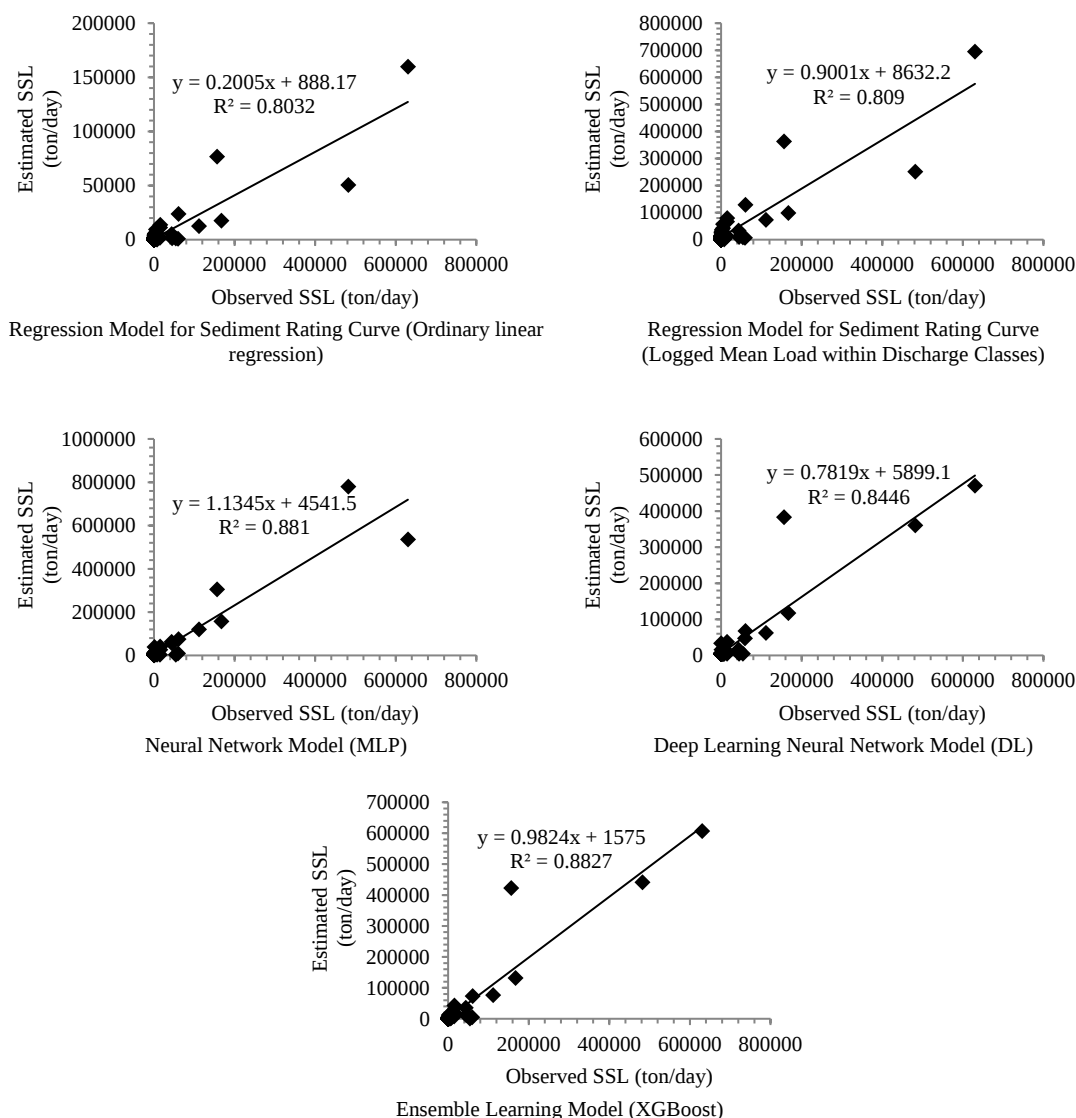
با توجه به نتایج مندرج در جدول ۵، مدل یادگیری جمعی (XGBoost) با داشتن کمترین میزان خطا MAE برابر با ۱۰۶۵۲/۱۳ تن در روز، RMSE برابر

۳۶۲۱۹ تن در روز و بیشترین شاخص کارایی (NSE) برابر با ۰/۸۷ به‌عنوان مدل منتخب از بین مدل‌های داده مینا برگزیده شد. در میان مدل‌های شبکه عصبی، مدل

منحنی سنج رسوب یک خطی است. از نقطه نظر میزان درصد بایاس نیز، مدل‌های یادگیر جمعی و عمیق به ترتیب با دارا بودن مقادیر ۳/۵۲- درصد بیش برآوردی و ۲/۰۴ درصد کم برآوردی، بهترین مدل‌های داده مینا محسوب می‌شوند.

در مجموع، با توجه به تحلیل آماری نتایج به‌دست آمده از جدول ۵، از مدل یادگیری جمعی (XGBoost) به‌منظور شبیه‌سازی دراز مدت رسوب معلق رودخانه اترک در بازه زمانی ۱۳۹۰ تا ۱۴۰۱ استفاده شد که نتایج آن در جدول ۶ و شکل‌های ۸ و ۹ آمده است.

شبکه عصبی یادگیری عمیق (DL) با داشتن میزان خطا MAE برابر با ۱۵۹۷۹/۰۲ تن در روز، RMSE برابر ۴۰۳۸۷/۱۲ تن در روز و شاخص کارایی (NSE) برابر با ۰/۸۴ و پس از آن مدل شبکه عصبی (MLP) با داشتن میزان خطا MAE برابر با ۱۶۰۱۵/۵۳ تن در روز، RMSE برابر ۴۵۰۵۲/۸۳ تن در روز و شاخص کارایی (NSE) برابر با ۰/۸ در رتبه‌های بعدی قرار دارند. در میان مدل‌های رگرسیونی نیز مدل منحنی سنج حد وسط دسته‌ها با داشتن میزان خطا MAE برابر با ۲۰۹۰۹/۲۳ تن در روز، RMSE برابر ۴۵۶۳۲/۰۸ تن در روز و شاخص کارایی NSE برابر با ۰/۸۱ بهتر از مدل



شکل ۷- نمودار پراکنش داده‌های مشاهداتی آزمون در مقابل داده‌های تخمینی در مدل‌های مختلف داده مینا

Fig. 7. Scatter plot of observed test data vs. estimated data in various data-driven models

جدول ۶- برآورد بار رسوب معلق (تن در ماه) دراز مدت ایستگاه هیدرومتری هوتن (۱۳۶۰-۱۴۰۱)

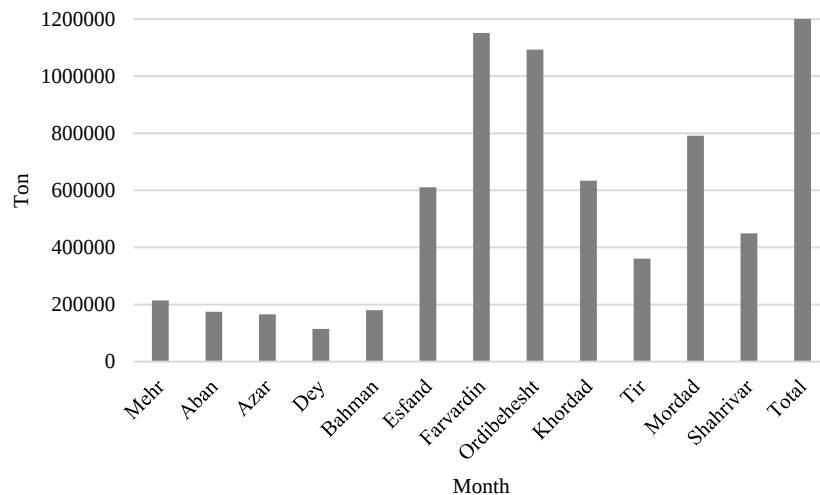
Table 6. Long-term estimation of suspended sediment load (tons per month) at Hootan Hydrometric Station (1981-2022)

Year	Mehr	Aban	Azar	Dey	Bahman	Esfand	Farvardin	Ordibehesht	Khordad	Tir	Mordad	Shahrivar	Total
1981	383220.5	77772.4	352931	254537.4	494634.1	807186.2	6958455	10374472	894792.2	2210348	1240172	18901.12	24067421
1982	1027622	286394.2	234753.3	206487.1	249799.5	166451	893898.9	38537.47	266625.2	239045	10566.27	1163334	4783513
1983	10846	29575.84	146311.4	26589.26	87349.88	87621.16	203706.6	33080.53	1355260	0	0	0	1980341
1984	101971.4	588406.3	328017.2	136681.1	340893.8	412766.5	696684.6	1171876	46502.36	770.91	3147.17	13650.16	3841367
1985	219105.9	28296.55	655608.7	140580	130606.7	37793.04	573034.6	135566.7	19376.7	0	4170646	8539.98	6119155
1986	50476.8	96127.17	1161807	97932.68	184491.5	1205847	572858.4	192857.4	147855.5	1078253	4679261	1024891	10492660
1987	343045.2	71443.33	46774.26	160503.2	319326.4	6525522	3914570	195581.3	875062.2	0	3081954	12.03	15533794
1988	30759.78	53728.98	291330.5	193851.9	283586.9	1155379	1812680	2910287	20798.82	330260.8	80610.34	2033715	9196989
1989	16833.6	25908.35	482408.9	125535.3	45962.55	100802.6	1567613	71559.36	52234.82	112114	118364.2	35745.56	2755082
1990	8082.52	18183.22	10489.56	24513.51	100165.6	180115.5	47368.31	515.4	2574.8	1441323	23874.01	0	1857205
1991	17772.04	14445.85	256827.8	133732.3	42051.82	235013.7	414930.8	205856.3	85423.91	2121.79	0	615651.1	2023827
1992	91292.57	132082.3	260978.1	331450.7	1537940	5798192	4949155	10522866	9180039	65406.67	21106.72	32171.17	32922680
1993	73310.72	212638	399443	759226.8	1145978	1193829	4304766	3449758	1292360	11042.77	1536165	66747.29	14445264
1994	251232.6	71004.48	245201.1	110965.8	73476.81	56291.78	622184.3	4692.91	50128.89	1835.2	0	1505355	2992369
1995	364957.9	25270.91	215315	22499.91	46908.05	257334.6	296065.6	189510.9	15.62	296797.3	0	7555.87	1722232
1996	133192.1	229910.9	20816.42	10217.81	33127.73	34363.68	435799.2	60087.53	417763.3	1304633	0	22724.28	2702636
1997	296460.2	1963338	100727.6	205225.8	198496.8	942876.1	65638.63	608115.2	4228737	406018.1	682.44	0	9016316
1998	25332.5	22284.78	23993.63	22316.47	31783.78	43918.56	2785943	1409648	663565.7	70175.26	1468983	47895.99	6615841
1999	1283571	992437	216478.5	24521.62	74890.06	72677.91	80830.53	1398.18	0	99736.12	4235969	0	7082510
2000	659089.3	15163.77	142218.9	18522.54	6982.5	15700.08	9953.26	58.32	119.02	19096.74	2285357	6337093	9509355
2001	150935	215980.8	10386.24	403309	217876.4	7648.21	222023.6	381705.4	52812.31	619127.1	2022830	18.84	4304653
2002	46.44	160078.7	22445.01	30945.11	101285	687219.1	1848650	3314521	23566.62	861629.6	1653167	246644.8	8950198
2003	71785.86	217924.4	118728.7	32198.07	93291.44	32391.2	2795780	2174709	2431936	20949.05	1764.08	30618.08	8022076
2004	63058.3	218555.2	502572.1	177696.5	513999	1015057	506195.2	134524.9	6707.16	49919.24	650248.7	647830	4486363
2005	134940	198112.1	205819.7	163970.4	74416.48	17777.33	675339.4	134257.5	148684.8	9041.79	2110913	11516.39	3884789
2006	9574.61	352887	77543.51	24162.45	14814.34	60451.74	21363.38	4.73	3142.59	0	0	24897.05	588841.4
2007	829936.3	7861.44	87222.79	28513.01	28860.43	29138.62	650695.1	23635.94	568646.9	5865.32	10022.42	15814.49	2286213
2008	5974.03	57.36	5735.22	1932.83	77844.3	790060.7	882.15	1415.03	37.24	1396.74	3.37	128692.2	1014031
2009	26408.15	4556.79	35825.13	18066.7	13533.87	315482.6	938176	967292.3	175186.5	3249.77	27.85	35.22	2497841
2010	12385.08	4032.2	388.36	2157.3	23240.83	89127.03	1730.37	125847.8	14834.53	750665.2	4563.84	2231995	3260968
2011	87003.96	19921.88	10806.57	9371.68	46591.2	111557.8	14817.35	92666.69	2408.46	1424.05	1881086	538115.5	2815771
2012	720114.2	750403.3	44029.07	73551.8	246504.7	48214.47	381334	477938	2009119	2289127	186787.5	90623.36	7317746
2013	369339.5	5198.82	43775.89	25427.53	4088.14	109.53	41955.23	545175.2	24961.76	13.64	15.13	35.52	1060096
2014	25.4	3398.49	48905.13	1347.94	19.86	24182.17	2202.11	26526.18	36498.17	0	388283.3	551480.4	1082869
2015	716075.3	108570.5	17517.16	313987.7	178129.3	14253.07	23238.1	16.99	4026.93	567192.9	406577.9	1175080	3524666
2016	66579.61	55023.37	8549.19	7560.16	9992.97	51762.76	188803.4	508737.8	116971.7	225774.5	25985.02	144120.8	1409861
2017	6170.42	0	4893.8	1658.87	8423.93	1230.77	17232.28	946.51	0.43	242521.8	94082.54	0	377161.3
2018	329601.4	15797.67	49942.38	461713	384850.5	2393091	1.79	1347714	5996.34	0	22691.03	1.79	5011401
2019	18576.31	25327.42	39348.65	21776.47	65577.01	613150.6	8045065	2468919	1388977	99354.84	123.44	1295.18	12787490

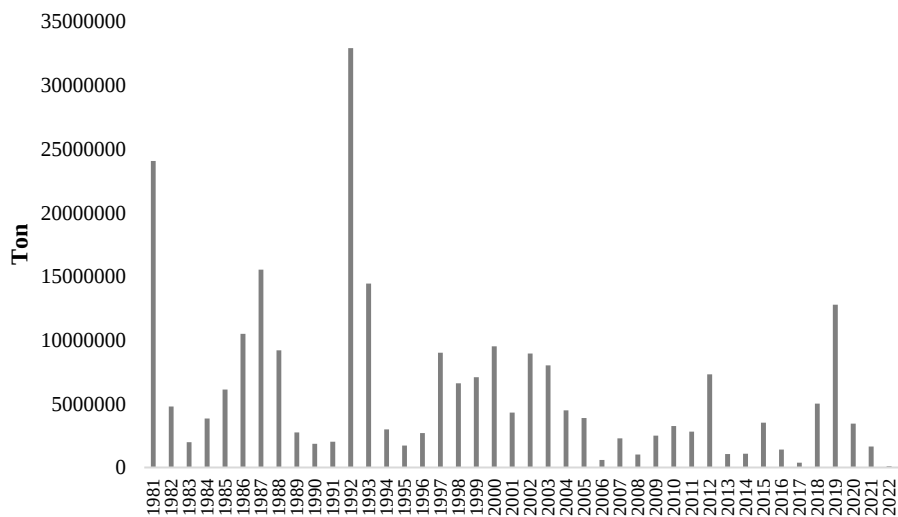
ادامه جدول ۶

Table 6 Continued

Year	Mehr	Aban	Azar	Dey	Bahman	Esfand	Farvardin	Ordibehesht	Khordad	Tir	Mordad	Shahrivar	Total
2020	10302.44	23157.59	24851.67	22825.75	20846.97	13744.23	757273.8	1531040	11834.67	1022526	5570.42	4146.66	3448120
2021	2054.92	33.01	62.12	719.86	25064.62	17115.06	5285.78	821.69	20.69	699539	797938.1	92653.21	1641308
2022	0	0	0	0	0	0	1481.54	76356.64	141.8	0	20477.76	24.75	98482.49
Ave. of all years (1360-1401)	214739.6	174792.6	165518.6	114971	180421.5	611010.6	1151087	1093121	633946.3	360911.8	791429	449276.7	5941226



شکل ۸- نمودار میانگین بار رسوب معلق (تن در ماه) ماه‌های مختلف سال در بازه ۱۳۶۰-۱۴۰۱ ایستگاه هیدرومتری هوتن
Fig. 8. Graph of the average suspended sediment load (tons per month) for different months of the year in the period 1981-2022 at the Hootan Hydrometric Station



شکل ۹- نمودار مجموع بار رسوب معلق (تن در سال) در بازه ۱۳۶۰-۱۴۰۱ ایستگاه هیدرومتری هوتن
Fig. 9. Graph of the total suspended sediment load (tons per year) in the period 1981-2022 at the Hootan hydrometric station

نخست سال ۱۴۰۱ مقدار ۵۹۴۱۲۲۶ تن در سال برآورد شده است. همچنین بیشترین میانگین بار رسوب معلق ماهانه مربوط به فرودین ماه با مقدار ۱۱۵۱۰۸۷ تن و کمترین آن مربوط به ماه دی با ۱۱۴۹۷۱ تن برآورد

با توجه به نتایج به دست آمده از شبیه‌سازی بار رسوب معلق رودخانه اترک و در محل ایستگاه هیدرومتری هوتن (جدول ۴)، میانگین دبی رسوب معلق سالانه این رودخانه در بازه زمانی ۱۳۶۰ تا شش ماه

در منطقه چات گنبد، شرایط انتقال رسوب ناشی از فرسایش قهقریایی آبراه‌های جانبی در پادگانه‌های آبرفتی-آسی، وضعیت ناپایداری برای کاربری‌های مستقر در حاشیه رودخانه اترک ایجاد کرده است. لذا، آگاهی از میزان انتقال رسوب معلق رودخانه در این منطقه اهمیت زیادی دارد. تغییرات شدید در غلظت رسوبات، بار رسوب و دبی جریان آب نشان‌دهنده تأثیر متغیرهای اقلیمی و هیدرولوژیکی است که مدیریت بهینه منابع آب و رسوب را ضروری می‌سازد.

این پروژه به‌منظور مکانیابی و طراحی حوضچه رسوبگیر، مخازن آب و ایستگاه پمپاژ در این منطقه انجام می‌شود. بنابراین، جمع‌آوری اطلاعات دقیق از میزان رسوبدهی رودخانه اترک در محل ایستگاه هوتن (نزدیک‌ترین ایستگاه هیدرومتری به سایت چات گنبد) امری ضروری است تا بتوان به بهترین شکل ممکن به نیازهای منابع آب و مدیریت رسوب پاسخ داد. لذا با توجه به وجود خلاءهای آماری فراوان در داده‌های رسوب معلق، از مدل‌های داده مبنا در شبیه‌سازی پیوسته بار رسوب معلق (در مقیاس روزانه) استفاده شد. در این رابطه و به منظور برآورد هر چه صحیح‌تر مقدار رسوب معلق رودخانه اترک (در محدوده طرح) از مدل‌های داده مبنا شامل منحنی سنج رسوب (حد وسط دسته‌ها و یک خطی)، مدل‌های هوشمند نظیر شبکه‌های عصبی MLP، SOM، یادگیری عمیق و یادگیری جمعی استفاده شد.

شیوه شبیه‌سازی رسوب معلق در این پژوهش دارای نوآوری بوده به‌نحوی که ابتدا بر اساس مدل جنگل تصادفی از میان متغیرهای دبی جریان (و روزهای پیشین آن) و بارش روزانه (و روزهای پیشین آن) متغیرهای با اهمیت در پیش‌بینی رسوب معلق تعیین و سپس با استفاده از یک مدل خوشه‌بندی، کل داده‌ها مورد خوشه‌بندی قرار گرفت. در مرحله بعد به‌منظور فرایند آموزش (واسنجی)، ارزیابی متقاطع و آزمون مدل‌ها، ضمن نمونه‌گیری از هر خوشه، دو یا سه مجموعه داده همگن و مشابه (به نسبت‌های ۷۰، ۱۵ و ۱۵ درصد) ساخته شد. شایان ذکر است که به‌دلیل احتمال وجود چولگی بالا در داده‌های رسوب معلق ایستگاه هوتن (به‌دلیل وضعیت رودخانه اترک و اراضی به‌شدت فرسایش منطقه)، تمهیداتی نظیر به‌کارگیری

شده است. در میان سال‌های آماری مورد بررسی (به استثناء سال ۱۴۰۱) بیشترین مقدار کل رسوب معلق سالانه مربوط به سال ۱۳۷۱ با مقدار $32922679/54$ تن و کمترین آن مربوط به سال ۱۳۹۶ با مقدار $377161/3$ تن بوده است.

در این پژوهش، مدل پیشنهادی مبتنی بر XGBoost به‌عنوان یک روش یادگیری جمعی، عملکردی برتر نسبت به مدل‌های سنتی منحنی سنج رسوب، شبکه‌های عصبی (ساده و یادگیری عمیق) در شبیه‌سازی و برآورد رسوب معلق نشان داد، به‌ویژه در رودخانه پرسوب اترک با داده‌های پرت و چولگی که نیازمند مدیریت نویز و ناهنجاری هستند. این مدل با بهره‌گیری از تکنیک‌های تقویت گرادیان، دقت بالایی در پیش‌بینی بار رسوبی فراهم می‌کند و در شرایط پیچیده رودخانه‌های کوهستانی یا حوضه‌های با بارندگی شدید کارایی بیشتری دارد.

مطالعات اخیر نیز این برتری را تأیید کرده‌اند. برای مثال، (Bezak et al., 2025) از XGBoost برای پیش‌بینی رسوب معلق در رودخانه‌های تحت تأثیر تغییرات اقلیمی استفاده کردند و دقت بالاتری نسبت به روش‌های سنتی گزارش نمودند. (Noh et al., 2023a) از مدل‌های یادگیری ماشینی برای تخمین دقیق سهم بار رسوب معلق از کل بار رسوبی در رودخانه‌های کم‌عمق استفاده نمود. نتایج نشان داد که مدل XGBoost به‌عنوان بهترین و دقیق‌ترین مدل در مقایسه با روش‌های سنتی و دیگر مدل‌های یادگیری ماشینی عمل نموده است. مطالعات (Piraei et al 2023) و (Kwon et al., 2023) نیز نتایج مشابهی را در خصوص کارایی مدل XGBoost در برآورد بار رسوبی رودخانه‌ها ارائه نمودند.

نتیجه‌گیری

در مطالعات منابع آب، اطلاعات دقیق و به‌هنگام درباره رسوب معلق رودخانه‌ها و پایش مستمر آنها برای آگاهی از شرایط رسوبدهی حوزه‌های آبخیز و تغییرات در بستر و حاشیه رودخانه‌ها بسیار حائز اهمیت است. این اطلاعات به‌ویژه در طراحی بهینه سازه‌های آبی و حجم مخازن ذخیره آب کاربرد دارد.

دقیق و مستمر از ایستگاه‌های هیدرومتری هوتن انجام گیرد تا امکان تحلیل‌های جامع‌تر و دقیق‌تری فراهم شود. در نهایت، اجرای برنامه‌های مدیریتی و اجرایی به‌منظور کنترل فرسایش و بهینه‌سازی منابع آب در مناطق مستعد رسوبگذاری، ضروری به‌نظر می‌رسد تا از پیامدهای منفی این پدیده‌ها در آینده جلوگیری شود. در خاتمه نتایج این پژوهش و تحلیل آماری داده‌های رسوب معلق ایستگاه هیدرومتری هوتن (به‌عنوان یکی از ایستگاه‌های هیدرومتری پر رسوب کشور) می‌تواند، به‌عنوان یک کار الگویی، در شبیه‌سازی رسوب معلق دیگر رودخانه‌ها یاریگر باشد.

تشکر و قدردانی

پژوهش انجام شده با حمایت مالی و معنوی پژوهشکده حفاظت خاک و آبخیزداری انجام شده است که بدین وسیله از کلیه حمایت‌ها و عزیزی که ما را در این راه کمک نموده‌اند، تشکر می‌شود.

تعارض منافع

در این مقاله تضاد منافی وجود ندارد و این مساله مورد تایید همه نویسندگان است.

توابع هدف مختلف سعی در بهبود فرایند آموزش و واسنجی مدل‌ها شد. در نهایت، تمامی مدل‌های پنجگانه با داده‌های آزمون یکسان مورد صحت‌سنجی قرار گرفت و با توجه به شاخص‌های صحت‌سنجی، مدل یادگیری جمعی XGBoost به‌عنوان مدل برتر انتخاب شد. در مرحله بعد، با استفاده از مدل منتخب، نسبت به برآورد درازمدت رسوب معلق ایستگاه هیدرومتری هوتن (سال‌های ۱۳۶۰ تا ۱۴۰۱) و تکمیل خلاءهای آماری در بازه زمانی ذکر شده اقدام شد.

در مجموع در پژوهش حاضر، علیرغم تلاش فراوانی که در حوزه به‌کارگیری و استفاده از مدل‌های مختلف داده مبنا در شبیه‌سازی دقیق رسوب معلق رودخانه اترک صورت گرفت. با این حال، این پژوهش با محدودیت‌هایی نیز مواجه بوده است. یکی از محدودیت‌های اصلی، وجود خلاءهای آماری و عدم قطعیت‌های موجود در داده‌های رسوب معلق است که می‌تواند بر دقت پیش‌بینی‌ها تأثیر بگذارد. علاوه بر این، تأثیر متغیرهای غیرقابل پیش‌بینی نظیر تغییرات اقلیمی و فعالیت‌های انسانی نیز می‌تواند بر نتایج مدل‌ها تأثیرگذار باشد. لذا برای مطالعات آتی، پیشنهاد می‌شود که تحقیقات بیشتری در زمینه جمع‌آوری داده‌های

منابع مورد استفاده

- Adnan, R.M., Petroselli, A., Heddham, S., Santos, C.A.G., Kisi, O., 2022. Suspended sediment modeling using a heuristic regression method hybridized with kmeans clustering. *Sustain.* 13(9), 4648. <https://doi.org/10.3390/su13094648>
- Bari, S.H., Yokoo, Y., Leong, C., 2024. A brief review of recent global trends in suspended sediment estimation studies. *Hydrolog. Res. Letters* 18(2), 51–57. <https://doi.org/10.3178/hrll.18.51>
- Bezak, N., Lebar, K., Bai, Y., Rusjan, S., 2025. Using machine learning to predict suspended sediment transport under climate change. *Water Resour. Manage.* 39, 3311–3326. <https://doi.org/10.1007/s11269-025-03809-0>
- Chachan, L.J., Bahnam, B.S., 2023. Long short-term memory for predicting monthly suspended sediment load. *Tianjin Daxue Xuebao (Ziran Kexue yu Gongcheng Jishu Ban)*. J. Tianjin Uni. Sci. Technol. 56(05), 36-49. <https://doi.org/10.17605/OSF.IO/R9MA2>
- Chen, T., Guestrin, C., 2016, August. Xgboost: A scalable tree boosting system. In *Proceedings of the 22nd acm sigkdd. International Conference on Knowledge Discovery and Data Mining*, 785-794.
- Doroudi, S., Sharafati, A., Mohajeri, S.H., 2021. Estimation of daily suspended sediment load using a novel hybrid support vector regression model incorporated with observer-teacher-learner-based optimization method. *Complex*. 2021(1), 5540284 (in Persian).
- Ezzaouini, M.A., Mahé, G., Kacimi, I., El Bilali, A., Zerouali, A., Nafii, A., 2022. Predicting daily suspended sediment load using machine learning and NARX hydro-climatic inputs in semi-arid environment. *Water* 14(6), 862. <https://doi.org/10.3390/w14060862>
- Etka, 2024. Data and reports collection and analysis of existing information: site selection and design of a sedimentation pond and water storage reservoir on the Atrak River and pumping station. Etka Organization, Tehran, Iran, 94 pages (in Persian).
- Ferguson, R.I., 1987. Accuracy and precision of methods for estimating river loads. *Earth Surf. Process. Landform.* 12(1), 95-104.

- Ghanbari-Adivi, E., 2025. A new machine learning model for predicting suspended sediment load. *J. Hydraulic*. 19(4), 31-45.
- Goodfellow, I., Bengio, Y., Courville, A., 2016. Deep learning. MIT Press, 708 pages.
- Hornik, K., Stinchcombe, M., White, H., 1989. Multilayer feedforward networks are universal approximators. *Neural Net*. 2(5), 359-366.
- Hosseiny, H., Masteller, C.C., Dale, J.E., Phillips, C.B., 2023. Development of a machine learning model for river bed load. *Earth Surf. Dynamic*. 11, 681-693. <https://doi.org/10.5194/esurf-11-681-2023>
- Humphreys, K.M., Mays, D.C., Water quality modeling and machine learning applied to suspended sediment and land management data from the South Fork Clearwater River Basin, Idaho County, Idaho, USA. Available at SSRN: <https://ssrn.com/abstract=4887564> or <http://dx.doi.org/10.2139/ssrn.4887564>
- Jansson, M.B., 1996. Estimating a sediment rating curve of the Reventazon river at Palomo using logged mean loads within discharge classes. *J. Hydrol*. 183(3), 227-241.
- Jarbais, G., Harshavardhanan, P., 2025. Comparative analysis of machine learning models for daily suspended sediment concentration prediction in environmental monitoring. *Sadhana* 50(63). <https://doi.org/10.1007/s12046-025-02705-1>
- Jones, K.R., Berney, O., Carr, D.P., Barrett E.C., 1981. Arid zone hydrology for agricultural development. FAO Irrigation and Drainage Paper, 37, 271 pages.
- Kaufman, L., Rousseeuw, P.J., 2009. Finding groups in data: An introduction to cluster analysis (Vol. 344). John Wiley and Sons, New York, USA, 342 pages.
- Kaveh, K., Kaveh, H., Bui, M.D., Rutschmann, P., 2021. Long short-term memory for predicting daily suspended sediment concentration. *Engin. Comput*. 37, 2013-2027.
- Khosravi, K., Golkarian, A., Saco, P.M., Booij, M.J., Melesse, A.M., 2022. Model identification and accuracy for estimation of suspended sediment load. *Geocarto Int*. 37(27), 18520-18545. doi:10.1080/10106049.2022.2142964
- Khosravi, M., Ghoochani, S., Shabani, H., 2024. Deep learning-based modeling of daily suspended sediment concentration and discharge in esopus. In 2024 International Conference on Machine Learning and Applications (ICMLA), 882-887). IEEE.
- Kohonen, T., 2002. The self-organizing map. *Proceedings of the IEEE*, 78(9), 1464-1480.
- Kumar, D., Bhattacharya, R.K., Singh, A., Banerjee, D., 2022. Modelling of suspended sediment concentration using conventional and machine learning approaches, in the Upper Godavari basin, India. *ISH J. Hydraulic Engin*. 28(1), 81-92. <https://doi.org/10.1080/09715010.2020.1723122>
- Kwon, S., Noh, H., Seo, I.W., Park, Y.S., 2023. Effects of spectral variability due to sediment and bottom characteristics on remote sensing for suspended sediment in shallow rivers. *Sci. Total Environ*. 878, 163125. <https://doi.org/10.1016/j.scitotenv.2023.163125>
- Lund, J.W., Groten, J.T., Karwan, D.L., Babcock, C., 2022. Using machine learning to improve predictions and provide insight into fluvial sediment transport. *Hydrolog. Process*. 36(8), e14648. <https://doi.org/10.1002/hyp.14648>
- Mir, A.A., Patel, M., 2024. A comprehensive review on sediment transport, flow dynamics, and hazards in steep channels. *J. Water Manage. Model*. 32, C517. <https://doi.org/10.14796/JWMM.C517>
- Mohamadi, S., 2019. The suspended sediment load modeling by artificial neural networks, neural-fuzzy and rating curve in Hlilrood watershed. *Watershed Engin. Manage*. 11(2), 452-466 (in Persian).
- Mohammadi-Raigani, Z., Gholami, H., Mohamadi, M., 2025. Evaluating the performance of machine learning models for predicting suspended sediment load, case study: Taleghan watershed, Iran. *E.E.R*. 15(1), 83-104 (in Persian).
- Mousavi, S.M., Elmizadeh, H., Sakiani, M.A., Zoratipour, A., 2024. Estimation of river-suspended sediments using ANNs and ANFIS methods with modeling in MATLAB, case study: Dez River in Khuzestan. *Oceanograph. Fisheries Open Access J*. 17(2), 555958. <https://doi.org/10.19080/OFOAJ.2024.17.555958>
- Noh, H., Son, G., Kim, D., Park, Y.S., 2023a. A novel efficient method of estimating suspended-to-total sediment load fraction in natural rivers. *Water Resour. Res*. 59(10), e2022WR034401. <https://doi.org/10.1029/2022WR034401>
- Noh, H., Son, G., Kim, D., Park, Y.S., 2023b. A SVR based-pseudo modified einstein procedure incorporating H-ADCP model for real-time total sediment discharge monitoring. *KSCE J. Civil Environ. Engin. Res*. 43(3), 321-335. <https://doi.org/10.12652/Ksce.2023.43.3.0321>
- Piraei, R., Afzali, S.H., Niazkar, M., 2023. Assessment of XGBoost to estimate total sediment loads in rivers. *Water Resour. Manage*. 37, 5289-5306. <https://doi.org/10.1007/s11269-023-03603-y>
- Rahgoshay, M., Feiznia, S., Arian, M., Hashemi, S.A.A., 2019. Simulation of daily suspended sediment load using an improved model of support vector machine and genetic algorithms and particle swarm. *Arab. J. Geosci*. 12(9), 277.

- Sahoo, B.B., Sankalp, S., Kisi, O., 2023. A novel smoothing-based deep learning time-series approach for daily suspended sediment load prediction. *Water Resour. Manage.* 37(11), 4271-4292.
- Shadkani, S., Hemmatzadeh, Y., Pak, A., Abolfathi, S., 2025. Prediction of suspended sediment concentration in fluvial flows using novel hybrid deep learning model. *Int. J. Sediment Res.* 40(4), 573-587. <https://doi.org/10.1016/j.ijsrc.2025.02.004>
- Shaukat, N., Hashmi, A., Abid, M., Aslam, M.N., Hassan, S., Sarwar, M.K., Masood, A., Shahid, M.L.U.R., Zainab, A., Tariq, M.A.U.R., 2022. Sediment load forecasting of gobindsagar reservoir using machine learning techniques. *Front. Earth Sci.* 10, 1047290. <https://doi.org/10.3389/feart.2022.1047290>
- Sit, M., Demiray, B.Z., Xiang, Z., Ewing, G.J., Sermet, Y., Demir, I., 2020. A comprehensive review of deep learning applications in hydrology and water resources. *Water Sci. Technol.* 82(12), 2635-2670.
- Swami, S., Underwood, K.L., Wshah, S., Davis, W.D., Rizzo, D.M., 2023. Advancing deep learning techniques for turbidity forecasting in rivers and reservoirs. *Graduate College Dissertations and Theses.* 1787. <https://scholarworks.uvm.edu/graddis/1787>
- Tabatabaei, M., Salehpour Jam, A., Mosaffaie, J., 2023. Suspended sediment simulation using machine learning algorithms and CHIRPS satellite precipitation data with emphasis on data clustering and gamma test, case study: Ramyan Watershed, Golestan Province. *Watershed Engin. Manage.* 15(3), 328-350.
- Taşar, B., Üneş, F., Demirci, M., Güzel, H., Varçin, H., 2024. Suspended sediment estimation using machine learning methods. 2024 "Air and Water – Components of the Environment" Conference Proceedings, Cluj-Napoca, Romania, 105-114. https://doi.org/10.24193/AWC2024_10.
- Tayfur, G., 2012. *Soft computing in water resources engineering: Artificial neural networks, fuzzy logic and genetic algorithms.* WIT Press, Dorset, UK, 288 pages.
- Ulke, A., Tayfur, G., Ozku, S., 2009. Predicting suspended sediment loads and missing data for gediz river, Turkey. *J. Hydrol. Engin.* 14(9), 954-965.
- Üneş, F., Tasar, B., Varçin, H., 2025. Forecasting of suspended sediment in river using artificial intelligence methods. *Air and Water – Components of the Environment Conference Proceedings, Cluj-Napoca, Romania*, 101-107. https://doi.org/10.24193/AWC2025_09
- Zhang, C., Liu, Y., Chen, X., Gao, Y., 2022. Estimation of suspended sediment concentration in the yangtze main stream based on sentinel-2 MSI data. *Remote Sens.* 14(18), 4446.