



Comparative Evaluation of Machine Learning Algorithms for Groundwater Level Modeling and Prediction: A Case Study of the Shahrood Plain

Sina Khoshnevisan¹ , Samad Emamgholizadeh² , Mohammadreza Asli charanadabi³ ,
Seyedeh Fatemeh Khakzad¹

1. MS.c Student, Department of Water Resources and Environmental Engineering, Faculty of Civil Engineering, Shahrood University of Technology, Shahrood, Iran.
2. Professor, Department of Water Resources and Environmental Engineering, Faculty of Civil Engineering, Shahrood University of Technology, Shahrood, Iran.
3. Ph.D Student, Department of Water Resources and Environmental Engineering, Faculty of Civil Engineering, Shahrood University of Technology, Shahrood, Iran.

✉Corresponding Author: s_gholizadeh517@shahroodut.ac.ir

Submit Date
07 November 2025

Revise Date
07 December 2025

Accept Date
07 December 2025

Keywords:

*Shahrood Aquifer,
Boosting Models,
Sustainable Water Resources
Management,
Climatic and Hydrological
Factors,
Groundwater Level
Fluctuations.*

Extended abstract

Introduction

This study aims to develop a comprehensive, data-driven evaluation of groundwater level dynamics in the Shahrood and Bastam aquifer system by comparing the predictive performance of five machine learning algorithms—XGBoost, CatBoost, Decision Tree (DT), Support Vector Regression (SVR), and K-Nearest Neighbors (KNN). Motivated by the persistent groundwater decline caused by excessive extraction, agricultural expansion, and climatic variability, the research pursues three specific goals: (i) integrate climatic drivers (precipitation and temperature), human-induced factors (well and qanat abstraction), and agricultural return flow into a unified predictive framework; (ii) quantify and compare the accuracy of competing algorithms using MAE, RMSE, and correlation coefficient (r) under an 80–20 train–test split; and (iii) identify the most reliable modeling approach for supporting groundwater management, extraction regulation, climate-impact assessment, and sustainability planning in arid and semi-arid aquifer systems.

Cite this article: Khoshnevisan, S., Emamgholizadeh, S., Asli charanadabi, M.R. & Khakzad, S.F. (2025). Comparative Evaluation of Machine Learning Algorithms for Groundwater Level Modeling and Prediction: A Case Study of the Shahrood Plain. *Journal of Aquifer and Qanat*, 6 (2), 77-96. DOI: 10.22077/jaaq.2025.10441.1132.



Copyright: © 2025 by the authors. Licensee Journal of Aquifer and Qanat. This article is an open-access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

Through this comparative evaluation, the study seeks to highlight the strengths and limitations of each algorithm and deliver actionable insights for water-resource managers tasked with mitigating long-term aquifer depletion

Materials and Methods

A 15-year monthly dataset (2000–2014) was assembled using observations from six precipitation stations, three temperature stations, groundwater abstraction records from 591 operational wells and qanats, and water-table measurements from 33 piezometric wells across the plain. Spatial averages of climatic variables were computed using the Thiessen polygon method, ensuring representation of spatial heterogeneity. The input matrix combines climatic, hydrological, and anthropogenic variables known to influence groundwater fluctuations in the semiarid Shahrood aquifer. The five machine learning models were calibrated under a systematic grid-search procedure to determine optimal hyperparameters. Model performance was evaluated using MAE, RMSE, and r to capture both magnitude-based and pattern-based predictive capability. This design enables fair assessment of contrasting modeling philosophies: tree-based boosting (XGBoost, CatBoost), instance-based learning (KNN), margin-based regression (SVR), and recursive partitioning (DT) under identical data and validation conditions

Results and Discussion

Groundwater levels showed a continuous declining trend over the study period, dropping from 1326.48 m in 2000 to 1315.68 m in 2014, an overall decline of 10.8 m, averaging 0.77 m per year. This pattern reflects the combined influence of insufficient recharge, reduced precipitation, rising temperatures, and intensifying extraction for agriculture, domestic use, and industry.

The model comparison demonstrated the clear superiority of gradient boosting methods. CatBoost achieved the lowest prediction error (MAE = 1.4029 m; RMSE = 1.9484 m), while XGBoost produced the strongest linear agreement with observed water-table fluctuations ($r = 0.8185$). Both algorithms outperformed classical models by a substantial margin, reducing RMSE by approximately 25–35% relative to DT, SVR, and KNN. The Decision Tree model exhibited limited accuracy (RMSE = 2.779 m; $r = 0.6701$), reflecting its inability to generalize under nonlinear, multivariate interactions. SVR provided slightly better pattern reproduction ($r = 0.6903$) but higher errors (RMSE = 2.6995), suggesting difficulty in capturing nonlinearities and noise-driven variability. KNN performed the weakest (RMSE = 2.8617 m; $r = 0.5799$), likely due to high sensitivity to noisy, heterogeneous hydro-climatic data. Overall, boosting algorithms' ensemble structure and capacity to model complex nonlinearities allowed them to reproduce both long-term declining trends and short-term fluctuations more accurately than traditional learners. The performance gap confirms that aquifer systems characterized by strong climatic–anthropogenic coupling benefit from high-capacity, ensemble-based predictors.

Conclusion

The extended comparative analysis demonstrates that machine learning, particularly gradient boosting algorithms, provides a reliable, scalable framework for modeling groundwater level changes in arid and semi-arid regions. XGBoost and CatBoost consistently outperformed classical models, achieving lower errors and higher correlation with observed groundwater levels. Their predictive strength arises from their ability to capture nonlinear interactions among climatic variables, extraction rates, and return flows interactions that simpler methods fail to fully represent.

The findings emphasize that boosting models can be integrated into groundwater monitoring systems to inform extraction control, evaluate climate-change impacts, and support sustainable aquifer management. These models offer practical value for policymakers by enabling early detection of critical declines and providing operational decision support for regulating pumping and designing recharge interventions. Limitations include the absence of land-use, soil, and hydraulic-property data, and the inability to incorporate deep multi-temporal dependencies that recurrent networks (e.g., LSTM) could capture. Future research should explore hybrid and deep-learning architectures, incorporate uncertainty quantification, and expand datasets using remote-sensing inputs. Nevertheless, the main implication is clear: ensemble-based machine learning is a powerful, cost-effective tool for predicting groundwater dynamics and guiding sustainable water-resource planning in data-limited, climatically stressed aquifers.



مقایسه عملکرد الگوریتم‌های یادگیری ماشین در مدل‌سازی و پیش‌بینی سطح آب زیرزمینی: مطالعه موردی دشت شاهرود

سینا خوشنویسان^۱ | صمد امامقلی‌زاده^۲ | محمدرضا اصلی‌چرندابی^۳ | سیده فاطمه خانزاد^۱

۱. دانشجوی کارشناسی ارشد، گروه مهندسی آب و محیط زیست، دانشکده مهندسی عمران، دانشگاه صنعتی شاهرود، شاهرود، ایران.

۲. استاد، گروه مهندسی آب و محیط زیست، دانشکده مهندسی عمران، دانشگاه صنعتی شاهرود، شاهرود، ایران.

۳. دانشجوی دکتری، گروه مهندسی آب و محیط زیست، دانشکده مهندسی عمران، دانشگاه صنعتی شاهرود، شاهرود، ایران.

✉ نویسنده مسئول: S_gholizadeh517@shahroodut.ac.ir

چکیده

کاهش مداوم سطح آب زیرزمینی در مناطق خشک و نیمه‌خشک، از جمله دشت شاهرود و بسطام، به‌واسطه برداشت بی‌رویه، توسعه کشاورزی و صنعتی و تغییرات اقلیمی، به یکی از چالش‌های جدی مدیریت منابع آب تبدیل شده است. این افت پیوسته پیامدهایی نظیر کاهش ذخیره مخزن، فرونشست زمین، افت کیفیت آب و تهدید پایداری کشاورزی را در پی دارد. ازاین‌رو، پیش‌بینی دقیق سطح آب زیرزمینی به‌عنوان ابزار مهمی برای برنامه‌ریزی و سیاست‌گذاری مبتنی بر مدیریت پایدار ضروری است. هدف این پژوهش ارزیابی و مقایسه عملکرد پنج الگوریتم یادگیری ماشین شامل CatBoost، XGBoost، درخت تصمیم (DT)، ماشین بردار پشتیبان (SVR) و K-نزدیک‌ترین همسایه (KNN) در پیش‌بینی سطح آب زیرزمینی دشت شاهرود و بسطام طی دوره ۲۰۰۰ تا ۲۰۱۴ است. داده‌های ورودی شامل متغیرهای اقلیمی (بارش و دما)، برداشت از چاه‌ها و قنات‌ها و آب برگشتی کشاورزی بوده و مدل‌ها با نسبت ۸۰ درصد آموزش و ۲۰ درصد آزمون، با شاخص‌های MAE، RMSE و ضریب همبستگی (r) ارزیابی شدند. نتایج نشان داد مدل‌های تقویتی عملکرد برتری نسبت به مدل‌های کلاسیک دارند؛ به‌گونه‌ای که CatBoost با MAE برابر ۱.۴۰۳ متر و RMSE برابر ۱.۹۴۸ متر کمترین خطا و XGBoost با r برابر ۰.۸۱۸ بیشترین همبستگی را به دست آورد. در مقابل، مدل‌های SVR، DT و KNN به دلیل محدودیت در شناسایی روابط غیرخطی، دقت پایین‌تری داشتند. یافته‌ها نشان می‌دهد مدل‌های Boosting می‌توانند به‌عنوان ابزار کارآمدی برای پیش‌بینی تراز آب زیرزمینی و پشتیبانی از تصمیم‌سازی در کنترل برداشت، ارزیابی اثر تغییر اقلیم و برنامه‌ریزی مدیریت پایدار آبخوان مورد استفاده قرار گیرند.

تاریخ دریافت: ۱۴۰۴/۰۸/۱۷

تاریخ بازنگری: ۱۴۰۴/۰۹/۱۶

تاریخ پذیرش: ۱۴۰۴/۰۹/۱۶

کلیدواژه‌ها:

آبخوان شاهرود،

مدل‌های Boosting

مدیریت پایدار منابع آب،

عوامل اقلیمی و هیدرولوژیکی،

نوسانات تراز آب.

مقدمه

در دهه‌های اخیر، رشد سریع جمعیت، توسعه فعالیت‌های کشاورزی و صنعتی و تغییر الگوهای بارش همراه با نوسانات اقلیمی، فشار قابل توجهی بر منابع آب وارد کرده است. در نتیجه، وابستگی به منابع آب زیرزمینی افزایش یافته و این منابع به یکی از اصلی‌ترین تأمین‌کنندگان آب شرب، کشاورزی و صنعت تبدیل شده‌اند (Noori et al., 2023). در بسیاری از مناطق جهان، و به‌ویژه در ایران، کاهش منابع آب سطحی و شرایط اقلیمی خشک موجب برداشت بی‌رویه و افت محسوس سطح آب زیرزمینی شده است. ایران، به‌ویژه در بخش مرکزی فلات، به دلیل بارش اندک و تبخیر بالا، اتکا و آسیب‌پذیری بیشتری نسبت به منابع زیرزمینی دارد؛ مسئله‌ای که موجب افت قابل توجه سطح آب در بسیاری از آبخوان‌ها شده است (Amanambu et al., 2020). افت سطح ایستابی، علاوه بر تشدید کسری مخزن، پیامدهای زیست‌محیطی و اقتصادی متعددی از جمله نشست زمین، کاهش کیفیت آب، شوری خاک و افت بهره‌وری کشاورزی را به همراه دارد. تغییر اقلیم نیز با افزایش دما، تغییر الگوی بارش و افزایش رخدادهای حدی مانند خشکسالی، موجب تغییرات مهمی در چرخه تغذیه و تخلیه آبخوان‌ها می‌شود (Klove et al., 2014). در چنین شرایطی، پایداری منابع آب زیرزمینی نیازمند شناخت دقیق رفتار آبخوان و توانایی پیش‌بینی تغییرات سطح آب در پاسخ به عوامل اقلیمی و انسانی است. از آنجا که سطح آب زیرزمینی شاخصی کلیدی برای ارزیابی وضعیت آبخوان‌ها محسوب می‌شود، پیش‌بینی دقیق آن نقشی مهم در مدیریت منابع آب، طراحی طرح‌های تغذیه مصنوعی و جلوگیری از افت بیش از حد سطح ایستابی دارد (Rahman et al., 2020).

مدل‌سازی رفتار آبخوان‌ها معمولاً با بهره‌گیری از مدل‌های عددی مبتنی بر معادلات فیزیکی انجام می‌شود، اما این مدل‌ها به دلیل نیاز به داده‌های گسترده هیدروژئولوژیکی،

زمان محاسباتی بالا و عدم قطعیت پارامترها، در بسیاری از کاربردها با محدودیت مواجه‌اند (Zhao et al., 2011). علاوه بر این، نوسان و ناهمگنی داده‌های ورودی در مقیاس‌های زمانی و مکانی مختلف، دقت مدل‌های سنتی را کاهش می‌دهد. به همین دلیل، در سال‌های اخیر توجه پژوهشگران به سمت روش‌های داده‌محور و به‌ویژه یادگیری ماشین جلب شده است (Hsu et al., 1995). روش‌های یادگیری ماشین با استفاده از داده‌های تاریخی، بدون نیاز به فرض روابط فیزیکی صریح، قادر به استخراج الگوهای پنهان و روابط پیچیده میان متغیرهای ورودی و خروجی هستند و با بهره‌گیری از الگوریتم‌های بهینه‌سازی، پیش‌بینی‌های دقیق‌تری ارائه می‌دهند (Da Silva et al., 2016). از مزیت‌های مهم این روش‌ها، انعطاف‌پذیری در مواجهه با داده‌های غیرخطی و ناقص و امکان به‌روزرسانی مداوم مدل‌ها با ورود داده‌های جدید است. تاکنون الگوریتم‌هایی مانند SVR، KNN، Decision Tree، XGBoost و CatBoost در مدل‌سازی پدیده‌های هیدرولوژیکی به‌کار رفته‌اند و هر یک با سازوکارهای متفاوت، توانایی مناسبی در شناسایی الگوهای غیرخطی و کاهش خطا از خود نشان داده‌اند. به‌عنوان نمونه، SVR قادر است روابط پیچیده را در فضاها با ابعاد بالا مدل‌سازی کند، در حالی که KNN بر شناسایی الگوهای محلی مبتنی است.

مدل‌های XGBoost و CatBoost به‌عنوان روش‌های تقویتی، با ترکیب تدریجی درخت‌ها خطای پیش‌بینی را کاهش می‌دهند و توانایی بالایی در شناسایی روابط غیرخطی دارند. با این حال، باید توجه داشت که این مدل‌ها - مانند سایر روش‌های داده‌محور - برای عملکرد مطلوب به داده‌های کافی و معتبر نیازمندند و در مناطق فاقد داده یا دارای کمبود شدید داده به‌تنهایی قابل استفاده نیستند. در چنین شرایطی، استفاده از داده‌های تکمیلی، درون‌یابی مقادیر گمشده یا بهره‌گیری از الگوهای

مناطق مشابه تنها در صورت وجود شرایط هیدروژئولوژیکی مشترک امکان‌پذیر است. بنابراین، کارایی این مدل‌ها زمانی بیشینه است که داده حداقلی لازم فراهم باشد.

بررسی پژوهش‌های گذشته نشان می‌دهد که استفاده از روش‌های یادگیری ماشین در مدل‌سازی سطح آب زیرزمینی طی دهه اخیر گسترش قابل‌توجهی داشته است. در یک مرور نظام‌مند شامل حدود ۲۰۰ مطالعه، مشخص شد که مدل‌های شبکه عصبی مصنوعی (ANN) بیشترین کاربرد را داشته و دقت آن‌ها نسبت به روش‌های خطی و رگرسیونی بالاتر است؛ همچنین کیفیت داده‌های ورودی، طول دوره مشاهدات و انتخاب متغیرهای کلیدی همچون بارش، دما و تبخیر بیشترین اثر را بر عملکرد مدل‌ها دارند (Ahmadi et al., 2022). در مطالعه‌ای دیگر در مناطق خشک ایران، مدل‌های RF، SVM و XGBoost برای پیش‌بینی تراز آب زیرزمینی به کار رفتند و نتایج حاکی از برتری XGBoost با ضریب تعیین بالاتر و خطای کمتر بود؛ بارش، تبخیر و برداشت انسانی نیز به عنوان عوامل اصلی مؤثر معرفی شدند (Feng et al., 2024). پژوهشی دیگر که چندین مدل داده‌محور از جمله رگرسیون چندمتغیره، RF و ANN را مقایسه کرد، نشان داد مدل SVR با کمترین RMSE و MAE بهترین عملکرد را دارد و افزودن ویژگی‌های مکانی مشتق‌شده از مدل مخلوط گاوسی موجب افزایش دقت پیش‌بینی می‌شود (Hussein et al., 2020). مرور فراتحلیلی ۱۹۷ مطالعه نیز تأکید کرد که ANN همچنان یکی از دقیق‌ترین مدل‌هاست و طول سری‌های زمانی (حداقل ۱۰ تا ۱۲ سال داده ماهانه) نقش مهم‌تری نسبت به نوع الگوریتم دارد (Ahmadi et al., 2022). مطالعه‌ای با مرور ۱۵۵ مقاله نشان داد که دما، بارش و سابقه نوسانات سطح آب مهم‌ترین ورودی‌ها هستند و ترکیب داده‌های چندمنبعی سبب افزایش قابل‌توجه دقت مدل‌ها می‌شود (El Ansari

et al., 2023). علاوه بر این، پژوهش‌های جدیدتر رویکردهای پیشرفته‌تری مانند AutoML-GWL را معرفی کرده‌اند که با بهینه‌سازی بیزی، تنظیم خودکار ابرپارامترها و ترکیب ۱۶ الگوریتم، دقت بالاتری ارائه می‌دهد (Singh et al., 2024). چارچوب چندمدلی ترکیبی بیزی نیز با استفاده از رویکرد Stacking توانسته است عدم قطعیت را کاهش دهد و پیش‌بینی‌های قابل اتکاتری ارائه کند (Zhu et al., 2025). در برخی مطالعات، مدل‌های ترکیبی مانند LSTM-SVR-ANFIS یا سیستم‌های فازی-عصبی نیز در پیش‌بینی نوسانات سطح آب به‌ویژه در شرایط اقلیمی متغیر عملکرد بهتری نسبت به مدل‌های منفرد نشان داده‌اند، (Emamgholizadeh, et al., 2014b). علاوه بر این، کاربرد یادگیری ماشین تنها به پیش‌بینی سطح آب محدود نشده و در حوزه‌هایی مانند مدل‌سازی کیفیت آب و تخمین تبخیر نیز نتایج موفقیت‌آمیزی گزارش شده است؛ به‌گونه‌ای که مدل‌های ANN، SVR، RF و روش‌های تقویتی توانسته‌اند روابط غیرخطی میان متغیرهای اقلیمی و پارامترهای کیفی آب یا تبخیر را با دقت بالایی شناسایی کنند.

بر اساس مطالعات گذشته، می‌توان نتیجه گرفت که استفاده از مدل‌های ترکیبی و خودکار (Ensemble و AutoML) در پیش‌بینی تراز آب زیرزمینی عملکرد بسیار بهتری نسبت به مدل‌های کلاسیک دارد. در بین الگوریتم‌های منفرد، LSTM و SVR به دلیل توانایی درک روابط غیرخطی و وابستگی‌های زمانی، دقت بالاتری نشان داده‌اند. همچنین، کیفیت داده‌های ورودی، انتخاب ویژگی‌های کلیدی و تحلیل عدم قطعیت بیزی از عوامل مؤثر بر کارایی نهایی مدل‌ها محسوب می‌شوند. ترکیب داده‌های اقلیمی، هیدروژئولوژیکی و انسانی در کنار الگوریتم‌های بهینه‌سازی خودکار، می‌تواند آینده‌ای دقیق‌تر و کارآمدتر را در مدیریت پایدار منابع آب رقم بزند (Rezapour & Sabzekar, 2025).

که الگوریتم‌های فرا-ابتکاری می‌توانند در تنظیم بهینه وزن‌ها و پارامترهای مدل‌های یادگیری ماشین نقش مهمی در بهبود دقت پیش‌بینی سطح آب زیرزمینی ایفا کنند (Nohani et al., 2025).

بر اساس مطالعات گذشته می‌توان نتیجه گرفت که مدل‌های ترکیبی و خودکار مانند Ensemble و AutoML در پیش‌بینی تراز آب زیرزمینی عملکرد دقیق‌تری نسبت به مدل‌های کلاسیک ارائه می‌دهند. در میان مدل‌های منفرد نیز الگوریتم‌هایی مانند LSTM و SVR به دلیل توانایی در شناسایی وابستگی‌های زمانی و روابط غیرخطی موفقیت بیشتری داشته‌اند. این یافته‌ها نشان می‌دهد که کیفیت داده‌های ورودی، انتخاب ویژگی‌های کلیدی و تحلیل عدم قطعیت از عوامل تعیین‌کننده در کارایی مدل‌های یادگیری ماشین هستند و به‌کارگیری داده‌های چندمنبعی (اقلیمی، هیدروژئولوژیکی و انسانی) همراه با روش‌های بهینه‌سازی می‌تواند دقت پیش‌بینی را به‌طور قابل توجهی افزایش دهد (Emamgholizadeh et al., 2014a).

اهمیت این موضوع در دشت شاهرود که یکی از مناطق کلیدی استان سمنان در بخش کشاورزی و صنعت است، دوچندان می‌شود. این آبخوان طی سال‌های اخیر با افت مستمر سطح آب مواجه بوده و عواملی مانند کاهش بارش، افزایش دما، برداشت بی‌رویه و کاهش تغذیه طبیعی آبخوان در تشدید این روند نقش داشته‌اند. پیامدهایی همچون کاهش ذخیره مخزن، فرونشست زمین، افت کیفیت منابع آبی و تهدید پایداری فعالیت‌های کشاورزی و صنعتی ضرورت استفاده از رویکردهای دقیق پیش‌بینی را برجسته می‌سازد. بنابراین، با توجه به تأکید مطالعات پیشین بر برتری مدل‌های پیشرفته، بهره‌گیری از الگوریتم‌های ترکیبی و داده‌محور برای تحلیل رفتار آبخوان شاهرود می‌تواند مبنایی علمی برای مدیریت پایدار این منبع حیاتی فراهم کند. (Ebrahimi et al., 2025).

در یک پژوهش به مقایسه عملکرد چند الگوریتم یادگیری ماشین شامل درخت تصمیم (DT)، جنگل تصادفی (RF) و ماشین بردار پشتیبان (SVM) در محیط GIS برای پیش‌بینی نوسانات سطح آب زیرزمینی دشت بیرجند پرداخته است. نتایج نشان داد مدل RF با دقت بالاتر (همبستگی ۰/۸۹) نسبت به سایر مدل‌ها توانایی بهتری در پیش‌بینی تغییرات مکانی و زمانی سطح آب زیرزمینی دارد. پژوهش تأکید می‌کند که استفاده از داده‌های اقلیمی و توپوگرافی به همراه شاخص‌های هیدروژئولوژیکی در ترکیب با الگوریتم‌های ML می‌تواند ابزار کارآمدی برای مدیریت منابع آب در مناطق خشک و نیمه‌خشک ایران باشد (Eftekhari et al., 2024).

در یک تحقیق از مدل‌های ترکیبی شبکه عصبی مصنوعی (ANN)، سیستم عصبی-فازی تطبیقی (ANFIS) و مدل موجک-عصبی (WNN) برای پیش‌بینی نوسانات سطح آب زیرزمینی در دشت جهرم استفاده شده است. نتایج حاکی از آن است که مدل ترکیبی WNN-ANFIS با در نظر گرفتن ویژگی‌های غیرخطی و نوسانی داده‌ها بهترین عملکرد را ارائه می‌دهد RMSE برابر با ۰/۲۳ متر نتیجه می‌گیرند که مدل‌های مبتنی بر یادگیری عمیق و داده‌های موجک، در مقایسه با مدل‌های کلاسیک، توانایی بیشتری در مدل‌سازی پویایی پیچیده سطح آب زیرزمینی دارند و برای پیش‌بینی کوتاه‌مدت و میان‌مدت قابل اطمینان‌تر هستند (Dastourani et al., 2025).

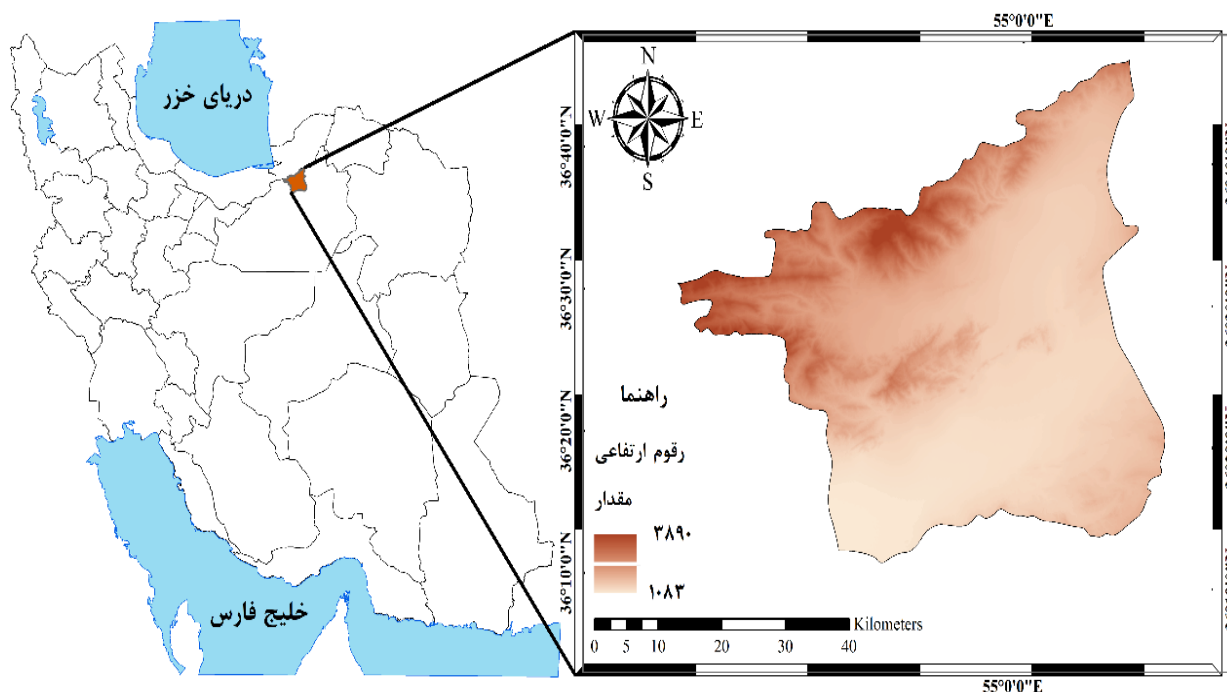
در یک مطالعه، کارایی مدل‌های فرا-ابتکاری شامل الگوریتم ازدحام ذرات (PSO) و الگوریتم ژنتیک (GA) در بهینه‌سازی مدل‌های هوش مصنوعی برای پیش‌بینی سطح آب زیرزمینی دشت دلفان بررسی شد. داده‌های هیدروژئولوژیکی و اقلیمی در بازه زمانی ۲۰ ساله برای آموزش مدل‌ها استفاده گردید. نتایج نشان داد ترکیب PSO-ANN نسبت به مدل‌های منفرد ANN یا GA دقت بالاتری دارد ($R^2=0.92$). مطالعه نتیجه گرفت

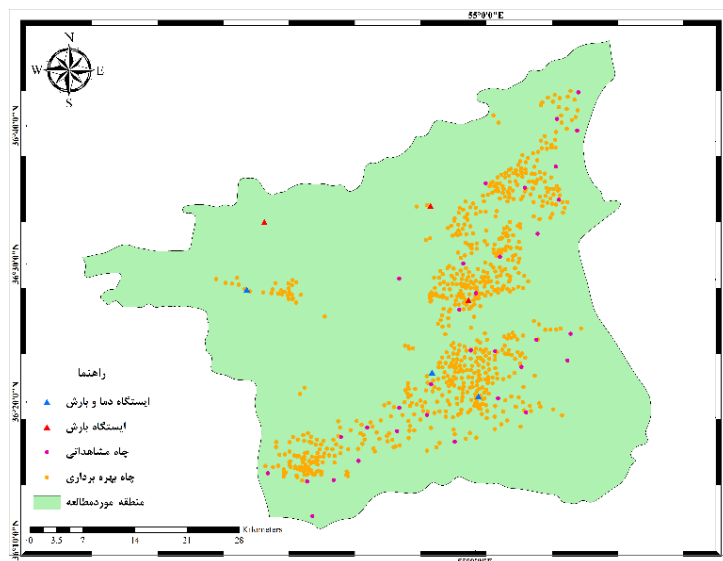
مواد و روش‌ها

منطقه مورد مطالعه

منطقه مورد مطالعه در بخش شمالی شهرستان شاهرود و در استان سمنان واقع شده است. این دشت از غرب به شهرستان دامغان، از شمال به استان گلستان، از شرق به خراسان شمالی و از جنوب به محدوده مرکزی شهرستان شاهرود محدود می‌شود. محدوده ارتفاعی منطقه بین حدود ۱۰۸۳ تا ۳۸۹۰ متر متغیر است که این تغییرات توپوگرافی بر الگوهای بارش و دما تأثیر می‌گذارد. اقلیم غالب منطقه نیمه‌خشک تا نیمه‌مرطوب بوده و میانگین بارش و دما در ایستگاه‌های مختلف آن تفاوت قابل توجهی دارد. منابع آب زیرزمینی مهم‌ترین منبع تأمین آب شرب، کشاورزی و صنعت منطقه محسوب می‌شوند و برداشت عمدتاً از طریق چاه‌ها و قنات‌ها انجام می‌گیرد. در شکل ۱، محدوده جغرافیایی دشت شاهرود همراه با موقعیت ایستگاه‌های باران‌سنجی و دماسنجی، چاه‌های پیزومتری و چاه‌های بهره‌برداری ارائه شده است.

بر این اساس، هدف پژوهش حاضر، بررسی و مقایسه‌ی عملکرد چند مدل یادگیری در پیش‌بینی سطح آب زیرزمینی دشت شاهرود است. در این مطالعه، داده‌های چاه‌های پیزومتری به همراه داده‌های اقلیمی، میزان برداشت آب از طریق چاه‌ها و قنات، و آب برگشتی کشاورزی برای دوره آماری ۲۰۰۰ تا ۲۰۱۴ به‌صورت ماهانه جمع‌آوری شده و برای آموزش و آزمون مدل‌ها مورد استفاده قرار گرفت. با ارزیابی دقت پیش‌بینی مدل‌ها بر اساس شاخص‌های مختلف ارزیابی، مدل مناسب برای پیش‌بینی سطح آب زیرزمینی مشخص گردید. نتایج این پژوهش می‌تواند در بهبود درک پویایی آبخوان شاهرود، تحلیل اثر تغییر اقلیم و برداشت بر رفتار آب زیرزمینی، و همچنین تدوین سیاست‌های مؤثر برای مدیریت پایدار منابع آب منطقه، می‌تواند راهنمای سیاست‌گذاری و مدیریت پایدار منابع آب باشد.

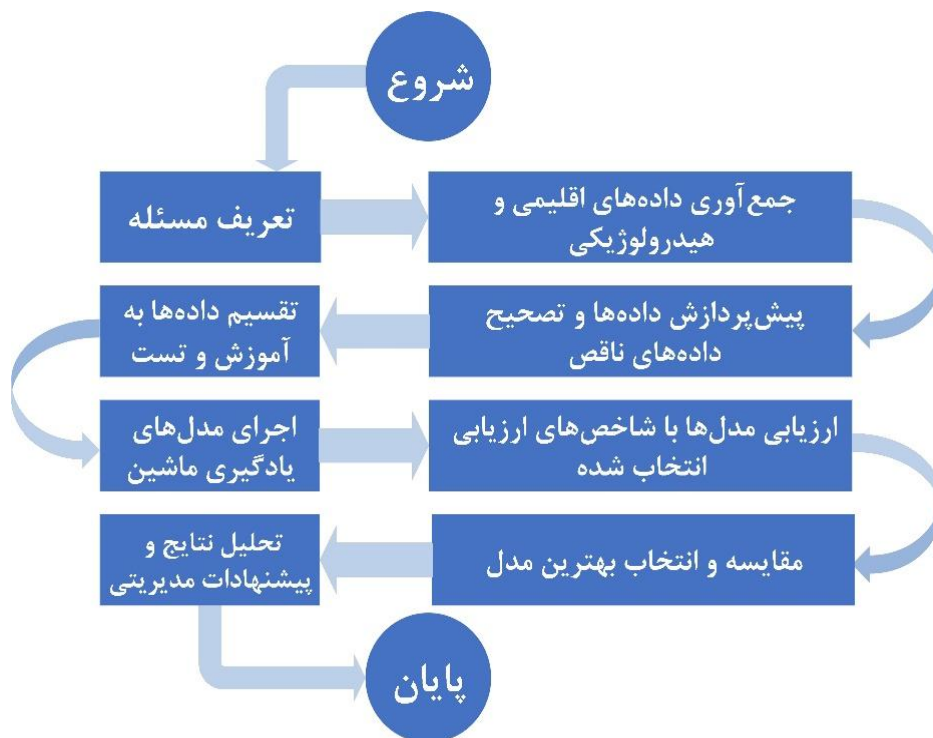




شکل ۱. محدوده دشت شهرستان شاهرود.
Fig 1. Extent of the Shahroud Plain.

مدل‌های مختلف یادگیری ماشین اجرا گردیدند. در ادامه، عملکرد مدل‌ها با شاخص‌های ارزیابی مقایسه شد و بهترین مدل انتخاب گردید. در نهایت، نتایج تحلیل و پیشنهادهای مدیریتی ارائه شدند.

شکل ۲ روند انجام پژوهش را از مرحله‌ی تعریف مسئله تا تحلیل نتایج نشان می‌دهد. ابتدا داده‌های اقلیمی و هیدرولوژیکی گردآوری و پیش‌پردازش شدند، سپس داده‌ها به دو بخش آموزش و آزمون تقسیم شدند و



شکل ۲. روندنما انجام پژوهش.
Fig 2. Research workflow.

داده‌های مورد استفاده

داده‌های بارش و دما

در پژوهش حاضر، داده‌های ماهانه شش ایستگاه باران‌سنجی واقع در منطقه مطالعه مورد استفاده قرار گرفت. ایستگاه سینوپتیک شاهرود تحت نظارت سازمان هواشناسی کشور است و پنج ایستگاه باقی‌مانده توسط وزارت نیرو اداره می‌شوند. برای متغیر دما نیز داده‌های سه ایستگاه به کار گرفته شد که شامل ایستگاه سینوپتیک شاهرود و دو ایستگاه تحت مدیریت وزارت نیرو می‌باشند. ایستگاه‌های انتخاب شده دارای دوره آماری کامل و پیوسته بیست‌ساله از سال ۲۰۰۰ تا ۲۰۱۴ هستند. برای

محاسبه میانگین فضایی بارش و دما در کل منطقه، از روش تیسن (Thiessen) استفاده شد. این روش با اختصاص وزن‌های متناسب به هر ایستگاه بر اساس مساحت حوزه تاثیر آن، امکان برآورد دقیق میانگین فضایی متغیرهای اقلیمی را فراهم می‌کند و پراکنش جغرافیایی ایستگاه‌ها را به‌طور کامل لحاظ می‌کند. مشخصات کامل ایستگاه‌ها در جدول ۱ ارائه شده است. در این مطالعه، ورودی‌های مدل شامل بارش، دما، برداشت از آب زیرزمینی از طریق چاه و قنات و نیز آب برگشتی کشاورزی بوده‌اند. این متغیرها نقش کلیدی در نوسانات سطح آب زیرزمینی دشت شاهرود و بسطام دارند.

جدول ۱. مختصات ایستگاه‌های بارش و دما پژوهش حاضر.

Table 1. Coordinates of Precipitation and Temperature Stations Used in The Present Study.

مؤلفه Parameter	نوع ایستگاه Station type	عرض جغرافیایی Latitude	طول جغرافیایی Longitude	کد ایستگاه Station code
بارش، دما Precipitation, Temperature	کلیماتولوژی Climatology	4039325	289487	47-033
بارش، دما Precipitation, Temperature	کلیماتولوژی Climatology	4025026	320566	47-046
بارش، دما Precipitation, Temperature	سینوپتیک Synoptic	4028198	314325	40-739
بارش Precipitation	کلیماتولوژی Climatology	4037910	319240	47-044
بارش Precipitation	کلیماتولوژی Climatology	4050537	314133	47-206
بارش Precipitation	کلیماتولوژی Climatology	4028198	291848	47-936

چاه‌های بهره‌برداری و پیزومتري

برای بررسی وضعیت منابع آب زیرزمینی در منطقه مطالعه، داده‌های مربوط به چاه‌های بهره‌برداری و چاه‌های پیزومتري مورد استفاده قرار گرفت. اطلاعات چاه‌های بهره‌برداری شامل ۵۹۱ چاه فعال است که برداشت آن‌ها در چهار گروه شرب، خدمات، صنعت و کشاورزی دسته‌بندی شده و بخشی از داده‌های ورودی مدل را تشکیل می‌دهد. همچنین، داده‌های تراز آب زیرزمینی از

۳۳ چاه پیزومتري طی دوره ۲۰۰۰ تا ۲۰۱۴ جمع‌آوری شد. داده‌های دارای نقص زمانی در این چاه‌ها با روش‌های مناسب تکمیل گردید تا پیوستگی سری زمانی حفظ شود. برای برآورد تراز متوسط آب زیرزمینی در سطح دشت نیز از روش تیسن استفاده شد. موقعیت نمونه‌ای از چاه‌های بهره‌برداری و تمامی چاه‌های پیزومتري در جداول ۲ و ۳ ارائه شده است.

جدول ۲. موقعیت نوع مصرف نمونه‌ای از چاه‌های بهره‌برداری.

Table 2. Location and Usage Type of A Sample of Exploitation Wells.

شماره Number	نام محدوده Region name	نام آبادی Settlement name	نوع مصرف Usage type	طول جغرافیایی Longitude	عرض جغرافیایی Latitude
1	بسطام	ابر	کشاورزی Agriculture	333469	4064772
2	بسطام	ابر	کشاورزی Agriculture	332550	4065207
3	شاهرود	حداده	شرب Drinking	297648	4017458
4	شاهرود	حداده	شرب Drinking	296622	4015517
5	شاهرود	مغان	صنعت Industry	317019	4027163
6	بسطام	گرمن	صنعت Industry	326387	4054996
7	شاهرود	خیرآباد	خدمات Services	331600	4034082
8	بسطام	ابر سج	کشاورزی Agriculture	317384	4048845
9	بسطام	پهن دره	شرب Drinking	314129	4037703
10	شاهرود	دیزج	صنعت Industry	324178	4024477

جدول ۳. مشخصات چاه‌های پیزومتری.

Table 3. Specifications of Piezometer Wells.

ایستگاه Station	طول جغرافیایی Longitude	عرض جغرافیایی Latitude	ایستگاه Station	طول جغرافیایی Longitude	عرض جغرافیایی Latitude
شرق کلاتخان East of Kalatkhan	309650	4020350	ابر جاده قطری Abr Jaddeh Qatri	333963	4065761
شرق مؤمن آباد East of Mowmenabad	292300	4014700	ابر گیاهی Abr Giyahi	331129	4062190
شرکت نفت Oil Company	322800	4031050	امیریه نکارمن Amiriyeh Nekarman	309932	4040815
شریف آباد Sharifabad	331374	4051371	پارک انقلاب Enghelab Park	319550	4031150
عباس آباد Abbasabad	301150	4013767	پری آباد Pariabad	298300	4008950
علی آباد Aliabad	304500	4016350	پلیس راه Police Station	314228	4026596
فرودگاه Airport	328350	4032600	تل Tal	317384	4018913
قلعه اصف Qaleh Asf	332500	4029800	جاده کلاتخان Kalateh-ye Khan Road	305644	4020834
قلعه بلوچ Qaleh Baluch	328503	4046793	جعفرخان Jafar Khan	326950	4022850
قلعه نو خالصه	313650	4022500	حصار	326350	4028950

Hesar			Qaleh Now-e Khalseh		
4023450	309950	خوریان Khourian	4019550	302150	کارگاه شن و ماسه Sand and Gravel Workshop
4033386	332936	خیر آباد Kheirabad	4043738	323455	کال چخماق Kal-e Chakhmaq
4060655	333822	دوراهی ابر اول جاده آب Dorah-e Abr Aval Jaddeh Ab	4052965	326789	کلامو Kolamu
4042793	318543	دولت آباد Dowlatabad	4036643	317997	مخزن Makhzan
4024740	323210	دیزج Dizaj	4013586	297596	مرادآباد Moradabad
4038825	320201	سیدابوالقاسم Seyed Abolghasem	4053584	321549	میغان Mighan
4055806	330955	نجفی Najafi			

شبکه‌های یادگیری ماشین

XGBoost (Extreme Gradient Boosting)

مدل XGBoost بر پایه‌ی ساختار درخت‌های تصمیم‌گیری تقویت‌شده طراحی شده است. این الگوریتم از یک معماری تجمعی و سلسله‌مراتبی بهره می‌برد که در آن مجموعه‌ای از درخت‌های تصمیم به صورت ترتیبی و افزایشی آموزش داده می‌شوند. در هر مرحله، مدل جدید با هدف اصلاح خطای باقیمانده‌ی مدل‌های پیشین ایجاد می‌شود و در نهایت ترکیب بهینه‌ای از درخت‌ها حاصل می‌گردد که توانایی بالایی در تقریب روابط غیرخطی میان متغیرها دارد. در ساختار داخلی XGBoost، هر درخت تصمیم به عنوان یک یادگیرنده‌ی ضعیف (Weak Learner) در نظر گرفته می‌شود که به طور تکراری و با وزن‌دهی مناسب در فرآیند به‌روزرسانی گرادینانی شرکت می‌کند. خروجی هر درخت، پیش‌بینی نسبی از مقدار هدف است و جمع وزنی این پیش‌بینی‌ها، نتیجه‌ی نهایی مدل را تشکیل می‌دهد. این فرآیند سبب می‌شود مدل به تدریج به سمت حداقل‌سازی تابع هزینه حرکت کرده و عملکرد آن در هر تکرار بهبود یابد. از دیدگاه محاسباتی، XGBoost از بهینه‌سازی مبتنی بر گرادینان دوم (Second-order

Gradient Optimization) بهره می‌گیرد؛ بدین معنا که علاوه بر گرادینانیکم، مشتقات دوم تابع هزینه نیز در فرآیند یادگیری لحاظ می‌شوند. این ویژگی موجب می‌شود مدل نسبت به تغییرات داده حساسیت کمتری داشته و از پایداری و همگرایی سریع‌تری برخوردار گردد. افزون بر این، الگوریتم از منظم‌سازی داخلی (Regularization) برای کنترل پیچیدگی مدل و جلوگیری از بیش‌برازش استفاده می‌کند. به کارگیری این منظم‌سازی‌ها باعث می‌شود شبکه از تعمیم‌پذیری بالاتری برخوردار باشد و بتواند داده‌های جدید را با دقت بیشتری پیش‌بینی کند. از منظر ساختاری، مدل شامل مجموعه‌ای از درخت‌های تصمیم با عمق‌های مختلف است که هر یک به صورت موازی ترتیبی در فرآیند آموزش مورد استفاده قرار می‌گیرند. این ساختار انعطاف‌پذیر و افزایشی به XGBoost امکان می‌دهد تا روابط پیچیده و غیرخطی میان متغیرهای ورودی و خروجی را مدل‌سازی کند؛ در حالی که سرعت آموزش آن به دلیل استفاده از روش‌های مبتنی بر بخش‌بندی هیستوگرامی (Histogram-based Binning) و محاسبه‌ی موازی بسیار بالاست.

CatBoost

مدل CatBoost Regressor بر پایه‌ی ساختار درخت‌های تصمیم‌گیری تقویت‌شده طراحی شده است. این الگوریتم همانند سایر مدل‌های مبتنی بر گرادیان تقویتی، از یک معماری تجمعی و افزایشی بهره می‌برد که در آن مجموعه‌ای از درخت‌های تصمیم به‌صورت ترتیبی و مرحله‌به‌مرحله آموزش داده می‌شوند. در هر گام، درخت جدیدی با هدف اصلاح خطای باقیمانده‌ی درخت‌های پیشین ساخته می‌شود و در نهایت، ترکیب بهینه‌ای از درخت‌ها حاصل می‌گردد که توانایی بالایی در مدل‌سازی روابط غیرخطی و پیچیده میان متغیرهای ورودی و خروجی دارد. در ساختار درونی CatBoost، هر درخت به‌عنوان یک یادگیرنده‌ی ضعیف (Weak Learner) ایفای نقش می‌کند که به‌صورت متوالی در فرآیند به‌روزرسانی گرادیانی شرکت دارد. خروجی هر درخت، پیش‌بینی نسبی از مقدار متغیر هدف است و مجموع وزنی این پیش‌بینی‌ها خروجی نهایی مدل را تشکیل می‌دهد. این فرایند سبب می‌شود مدل در هر تکرار، با کمینه‌سازی تدریجی تابع خطا، به سمت همگرایی بهینه حرکت کند و دقت پیش‌بینی افزایش یابد. ویژگی برجسته CatBoost نسبت به مدل‌های مشابه، استفاده از سازوکار Ordered Boosting است که با ترتیب‌دهی مناسب داده‌های آموزشی، از بروز پدیده‌ی بایاس پیش‌نگری (Prediction Shift) جلوگیری می‌کند. این روش موجب افزایش پایداری و تعمیم‌پذیری مدل در داده‌های واقعی می‌شود. علاوه بر این، CatBoost با به‌کارگیری تابع منظم‌سازی L_2 (L2 Regularization)، کنترل موثری بر پیچیدگی مدل و جلوگیری از بیش‌برازش اعمال می‌کند. از نظر محاسباتی، این مدل از مفاهیم بخش‌بندی هیستوگرامی (Histogram-based Binning) برای تسریع فرآیند تقسیم ویژگی‌ها و افزایش کارایی استفاده می‌کند. همچنین پارامترهایی نظیر عمق درخت (Depth)، نرخ یادگیری

(Learning Rate)، دمای نمونه‌گیری (Bagging Temperature) و تعداد مرزهای ویژگی (Border Count) نقش کلیدی در کنترل تعادل میان دقت و سرعت آموزش دارند.

KNN (K-Nearest Neighbors)

الگوریتم KNN یک روش یادگیری مبتنی بر نمونه است که برای مسائل رگرسیون کاربرد دارد و بر اساس شباهت بین نقاط داده عمل می‌کند. در این رویکرد، ابتدا پارامتر k که نشان‌دهنده‌ی تعداد همسایگان مورد بررسی برای پیش‌بینی است، تعیین می‌شود. سپس الگوریتم فاصله بین نمونه‌ی جدید و هر یک از نمونه‌های مجموعه‌ی آموزشی را محاسبه می‌کند، که معمولاً فاصله اقلیدسی (Euclidean Distance) به‌عنوان معیار فاصله به‌کار می‌رود. پس از شناسایی k همسایه‌ی نزدیک، مقدار پیش‌بینی شده برای نمونه‌ی جدید با میانگین‌گیری از مقادیر هدف این همسایگان به دست می‌آید. با توجه به حساسیت این الگوریتم نسبت به مقیاس ویژگی‌ها، استانداردسازی داده‌ها پیش از آموزش ضروری است تا همه‌ی ویژگی‌ها به‌طور مساوی در محاسبه‌ی فاصله مؤثر باشند و از ایجاد سوگیری ناشی از تفاوت دامنه مقادیر ویژگی‌ها جلوگیری شود. با توجه به ماهیت غیرپارامتریک و اتکای KNN بر آورد مبتنی بر نمونه‌های محلی، این الگوریتم به‌ویژه در تحلیل داده‌هایی با توزیع‌های پیچیده یا ناشناخته عملکرد مؤثری دارد و قادر است الگوهای غیرخطی و روابط محلی میان ویژگی‌ها و متغیر هدف را به‌خوبی شناسایی کند.

Decision Tree

الگوریتم درخت تصمیم (Decision Tree) یک روش یادگیری ماشین است که به‌طور گسترده در مسائل رگرسیون به‌کار می‌رود و بر اساس تقسیم‌بندی بازگشتی فضای ویژگی‌های ورودی عمل می‌کند. در این روش، داده‌ها به‌صورت مرحله‌به‌مرحله به زیرمجموعه‌های

کوچک‌تر تقسیم می‌شوند تا ساختاری درختی سلسله‌مراتبی شکل گیرد. در هر گره از درخت، یک ویژگی مشخص به‌عنوان معیار تصمیم‌گیری انتخاب می‌شود و داده‌ها بر اساس مقدار آستانه‌ای مرتبط با آن ویژگی به شاخه‌های مجزا تقسیم می‌شوند. انتخاب ویژگی بهینه و مقدار آستانه متناظر معمولاً بر اساس معیارهای از نوع خطا انجام می‌شود. این فرآیند بازگشتی تا رسیدن به شرایط توقف از پیش تعیین‌شده ادامه می‌یابد؛ به‌عنوان مثال، رسیدن به حداقل تعداد نمونه‌ها در یک گره یا دستیابی به سطح رضایت‌بخش کاهش واریانس. برای پیش‌بینی، مقدار خروجی برای یک نمونه‌ی جدید با میانگین‌گیری از مقادیر هدف نمونه‌های آموزشی موجود در گره برگ متناظر برآورد می‌شود. از ویژگی‌های برجسته‌ی درخت‌های تصمیم می‌توان به عدم حساسیت به مقیاس ویژگی‌ها اشاره کرد، که باعث می‌شود نیازی به استانداردسازی داده‌های ورودی نباشد. علاوه‌براین، با توجه به قابلیت تفسیر بالا و توانایی در شناسایی روابط غیرخطی میان ویژگی‌ها، درخت‌های تصمیم به‌طور گسترده در مسائل رگرسیون مختلف مورد استفاده قرار می‌گیرند.

SVR (Support Vector Regression)

الگوریتم SVR (Support Vector Regression) یک تکنیک یادگیری ماشین است که ریشه در نظریه‌ی ماشین بردار پشتیبان (SVM) دارد و برای مدل‌سازی روابط پیچیده و غیرخطی میان متغیرها در مسائل رگرسیون طراحی شده است. برخلاف روش‌های رگرسیون سنتی که به کمینه‌سازی خطای مطلق یا مربعات خطا می‌پردازند، SVR هدف خود را بر یافتن یک ابرصفحه بهینه (Hyperplane) قرار می‌دهد که مقادیر هدف را در محدوده‌ای مشخص از خطا، موسوم به اپسیلون (ϵ) پیش‌بینی کند. بدین ترتیب، الگوریتم قادر است با حفظ دقت مدل در چارچوب حاشیه‌ی خطای قابل قبول،

حساسیت به داده‌های پرت را کاهش دهد. برای شناسایی الگوهای غیرخطی، SVR داده‌های ورودی را با استفاده از توابع هسته‌ای (Kernel Functions) به فضای ویژگی‌های با بعد بالاتر نگاشت می‌دهد. این توابع می‌توانند شامل هسته‌های خطی، چندجمله‌ای (Polynomial) یا هسته تابع پایه شعاعی (RBF) باشند. در فضای نگاشت‌شده، الگوریتم دو ابرصفحه‌ی موازی تعریف می‌کند که بخش عمده‌ای از داده‌های آموزشی را در بر می‌گیرد. تنها نمونه‌هایی که خارج از این ناحیه‌ی حساس به ϵ قرار دارند، موسوم به بردارهای پشتیبان (Support Vectors)، تأثیر مستقیم بر شکل‌گیری مدل نهایی دارند. فرآیند بهینه‌سازی در SVR شامل کمینه‌سازی تابع هزینه است که انحرافات خارج از حاشیه ϵ را با استفاده از پارامتر منظم‌سازی C جریمه می‌کند. این پارامتر تعادلی میان پیچیدگی مدل و میزان تحمل خطا ایجاد می‌کند. تنظیم دقیق مقادیر C و ϵ برای دستیابی به مدلی با قابلیت تعمیم بالا و جلوگیری از بیش‌برازش و کم‌برازش بسیار حیاتی است.

جستجوی شبکه‌ای پارامترها (Grid Search)

به‌منظور بهینه‌سازی عملکرد مدل‌های پیش‌بینی و شناسایی ترکیب پارامترهای بهینه، در این پژوهش از جستجوی شبکه‌ای Grid Search استفاده شد. در این روش، فضای تعیین‌شده‌ای از مقادیر مختلف پارامترهای اصلی هر مدل تعریف گردید و تمامی ترکیب‌های ممکن به‌صورت سیستماتیک مورد ارزیابی قرار گرفتند. برای هر ترکیب پارامتر، مدل بر روی داده‌های آموزشی آموزش داده شد و عملکرد آن بر روی مجموعه داده‌های آزمون سنجیده شد. این فرآیند امکان شناسایی پارامترهایی که بیشترین دقت پیش‌بینی و کمترین خطای مدل را فراهم می‌کنند، فراهم آورد. استفاده از Grid Search، به‌ویژه برای مدل‌های پیچیده‌ای که شامل چندین پارامتر حساس

پایان دوره به ۱۳۱۵/۶۸ متر کاهش یافته است. این تغییر معادل افتی در حدود ۱۰/۸ متر طی چهارده سال و میانگین افت سالانه حدود ۰/۷۷ متر می‌باشد. این داده‌ها مربوط به بازه‌ی کل مطالعاتی پژوهش حاضر بوده و داده‌های آموزش و آزمون مدل‌ها را نیز شامل می‌شوند؛ همچنین با هدف نمایش روند کلی و تحلیل تغییرات سطح آب زیرزمینی در طول زمان ارائه شده‌اند. روند نزولی مشاهده شده در نمودار بیانگر آن است که نرخ برداشت از آبخوان به‌طور قابل ملاحظه‌ای از میزان تغذیه طبیعی آن فراتر رفته و در نتیجه منجر به کاهش مستمر ذخیره آب زیرزمینی شده است. در برخی سال‌ها، به‌ویژه بین سال‌های ۲۰۰۹ تا ۲۰۱۲، نوسانات جزئی در منحنی سطح آب مشاهده می‌شود که احتمالاً ناشی از تغییرات موقتی در میزان بارش یا محدودیت برداشت بوده است، اما این تغییرات کوتاه‌مدت نتوانسته‌اند روند کلی کاهش را متوقف کنند. افت مستمر سطح ایستابی را می‌توان حاصل تأثیر هم‌زمان عوامل انسانی و اقلیمی دانست؛ از یک‌سو برداشت بی‌رویه از چاه‌ها و توسعه‌ی بی‌برنامه‌ی اراضی کشاورزی موجب کاهش شدید منابع آب شده و از سوی دیگر، کاهش بارش مؤثر، افزایش دما و تشدید تبخیر ناشی از تغییر اقلیم، روند افت آب زیرزمینی را تشدید کرده است.

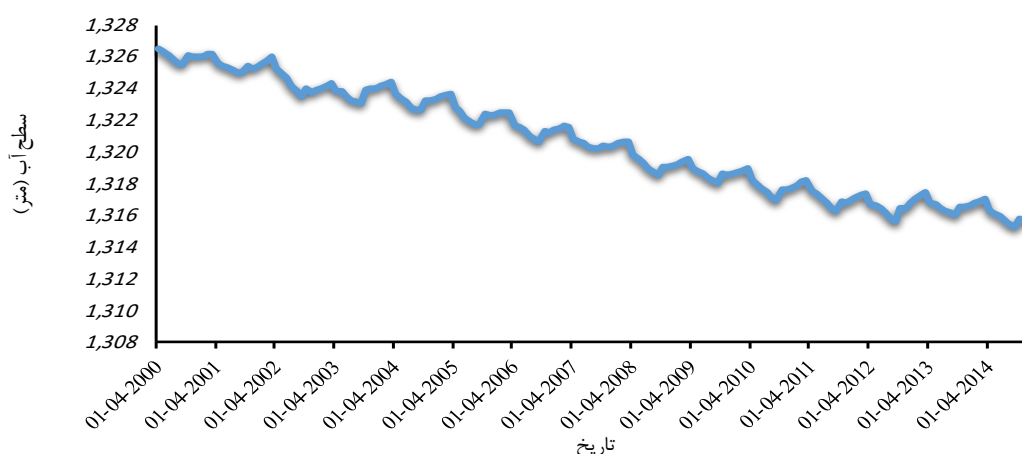
به عملکرد هستند، موجب شد تا از انتخاب مقادیر غیرمناسب جلوگیری شده و مدل‌ها با توان تعمیم بالاتر بر داده‌های جدید آموزش داده شوند. به این ترتیب، هر مدل با ترکیب بهینه پارامترهای خود ساخته شد و برای تحلیل دقیق‌تر روابط میان ویژگی‌های ورودی و خروجی آماده گردید. این روش، اطمینان حاصل کرد که مقایسه‌ی عملکرد مدل‌ها بر اساس بهترین حالت هر مدل انجام شود و ارزیابی‌ها قابل اتکا باشند.

شاخص‌های ارزیابی

برای ارزیابی عملکرد مدل‌ها از سه شاخص $RMSE$ ، MAE و ضریب همبستگی r استفاده شد. شاخص MAE میانگین خطای مطلق بین مقادیر واقعی و پیش‌بینی‌شده را نشان می‌دهد و $RMSE$ ریشه میانگین مربعات خطا است که میزان انحراف پیش‌بینی‌ها از مقادیر واقعی را بیان می‌کند. ضریب همبستگی r نیز میزان رابطه خطی میان مقادیر واقعی و پیش‌بینی‌شده را مشخص می‌سازد.

نتایج و بحث

نتایج حاصل از بررسی داده‌های مشاهداتی در شکل ۳، نشان می‌دهد که سطح آب زیرزمینی در دشت شاهرود طی دوره‌ی زمانی ۲۰۰۰ تا ۲۰۱۴ میلادی روندی نزولی و مداوم را تجربه کرده است. مقدار سطح آب زیرزمینی در ابتدای دوره برابر با ۱۳۲۶/۴۸ متر از سطح دریا بوده و تا



شکل ۳. سطح آب زیرزمینی در دشت شاهرود.

Fig 3. Groundwater Level in The Shahrud Plain.

دوره زمانی مورد استفاده (۲۰۰۰ تا ۲۰۱۴) نیز شامل سال‌هایی با تغییرات اقلیمی محسوس است که بر شدت برداشت و تغذیه آبخوان اثرگذار بوده است. به منظور پیش‌بینی سطح آب زیرزمینی از مدل‌های مختلف یادگیری ماشین استفاده گردید. جدول ۱ و شکل ۲ مقادیر شاخص‌های ارزیابی مدل‌ها شامل میانگین قدر مطلق خطا (MAE)، ریشه میانگین مربعات خطا (RMSE) و ضریب همبستگی (r) را نشان می‌دهد. این شاخص‌ها برای مدل‌های مختلف از جمله CatBoost، XGBoost، درخت تصمیم (DT)، K-نزدیک‌ترین همسایه (KNN) و ماشین بردار پشتیبان (SVR) محاسبه شده‌اند. بر اساس نتایج جدول فوق، مدل XGBoost با مقدار ضریب همبستگی برابر با ۰/۸۱۸، بالاترین ضریب همبستگی بین مقادیر پیش‌بینی شده و مشاهداتی را دارد، که نشان‌دهنده توانایی این مدل در بازتولید الگوهای تغییر سطح آب زیرزمینی است. مقدار شاخص ارزیابی MAE برابر با ۱/۴۸۹۷ و RMSE برابر با ۱/۹۶۹۲ نیز نشان می‌دهد که میانگین اختلاف مطلق بین مقادیر واقعی و پیش‌بینی شده کمتر از ۱/۵ متر است و خطای کلی مدل نیز کمتر از ۲ متر بوده است. مدل CatBoost نیز عملکردی بسیار نزدیک به XGBoost داشته است؛ MAE برابر با ۱/۴۰۲۹ و RMSE برابر با ۱/۹۴۸۴، مقدار خطای مطلق آن حتی اندکی کمتر از XGBoost است، اما به دلیل ضریب همبستگی کمتر ($r = 0.7815$)، دقت همبستگی زمانی پایین‌تری داشته است. این موضوع نشان می‌دهد که CatBoost در سطح خطای مطلق عملکرد خوبی دارد، اما در شناسایی روندهای زمانی تغییرات سطح آب کمی ضعیف‌تر است. مدل‌های درخت تصمیم (DT)، ماشین بردار پشتیبان (SVR) و KNN نسبت به دو مدل فوق عملکرد ضعیف‌تری نشان داده‌اند. برای مدل DT، مقادیر MAE برابر با ۱/۸۷۸۷ و RMSE برابر با ۲/۷۷۹۰ بوده که بیانگر افزایش خطا و کاهش دقت است. ضریب

همبستگی پایین‌تر آن ($r = 0.6701$) نیز نشان می‌دهد که این مدل قادر به درک کامل وابستگی‌های غیرخطی بین متغیرهای ورودی (مانند بارش، دما و برداشت از چاه و قنات) و سطح آب نبوده است. مدل SVR با MAE برابر با ۱/۹۵۸۴ و RMSE برابر با ۲/۶۹۹۵ عملکردی مشابه DT داشته ولی با ضریب همبستگی ۰/۶۹۰۳ اندکی بهتر بوده است. در نهایت، مدل KNN ضعیف‌ترین عملکرد را با MAE برابر با ۲/۲۵، RMSE برابر با ۲/۸۶۱۷ و r برابر با ۰/۵۷۹۹ از خود نشان داده است که نشان می‌دهد این الگوریتم در شناسایی روابط پیچیده و داده‌های غیرخطی منطقه مورد مطالعه مناسب نیست. به منظور تحلیل دقیق‌تر تفاوت عملکرد مدل‌ها، مقایسه‌ی نسبی بین شاخص‌های ارزیابی انجام شد. نتایج نشان داد که مدل XGBoost در مقایسه با درخت تصمیم (DT) توانست مقدار RMSE را حدود ۲۹ درصد کاهش دهد (از ۲/۷۷۹۲ به ۱/۹۶۹۲ متر)، که بیانگر بهبود قابل توجه در دقت کلی پیش‌بینی است. همچنین مقدار MAE در این مدل نسبت به DT حدود ۲۱ درصد کمتر بوده است (۱/۴۸۹۷ در برابر ۱/۸۷۸۷ متر). از سوی دیگر، ضریب همبستگی (r) در مدل XGBoost نسبت به مدل KNN حدود ۲۴ درصد افزایش داشته است (از ۰/۵۷۹۹ به ۰/۸۱۸۵)، که نشان‌دهنده توانایی بالای این مدل در بازنمایی الگوی واقعی تغییرات سطح آب زیرزمینی است. مدل CatBoost نیز در مقایسه با مدل‌های کلاسیک‌تر بهبود قابل توجهی نشان داده است، به طوری که مقدار RMSE آن حدود ۳۰ درصد کمتر از KNN و ۲۶ درصد کمتر از SVR است. این کاهش خطا همراه با مقدار نسبتاً بالای ضریب همبستگی ($r = 0.7815$) بیانگر پایداری و دقت بالای این مدل در پیش‌بینی داده‌های پیچیده‌ی هیدرولوژیکی است. در مجموع، مقایسه‌ی عددی شاخص‌ها نشان می‌دهد که مدل‌های مبتنی بر تقویت گرادیان (Boosting) نسبت به سایر الگوریتم‌های مورد استفاده، در حدود ۲۵ تا ۳۵

موقعیت ایده‌آل پیش‌بینی کامل را مشخص می‌کند. مطابق شکل، نقاط مربوط به مدل‌های CatBoost و XGBoost بیشترین نزدیکی را به خط برابری دارند، که نشان‌دهنده توانایی بالای این مدل‌ها در بازتولید الگوهای واقعی تغییر سطح آب زیرزمینی است. تراکم و هم‌گرایی بالای نقاط این دو مدل در اطراف خط برابری حاکی از آن است که آن‌ها توانسته‌اند روابط غیرخطی پیچیده میان متغیرهای ورودی (شامل بارش، دما، برداشت از چاه و قنات، و آب برگشتی کشاورزی) و سطح آب زیرزمینی را به خوبی شناسایی کنند. جدول ۴ عملکرد مدل‌های یادگیری ماشین را نشان می‌دهد.

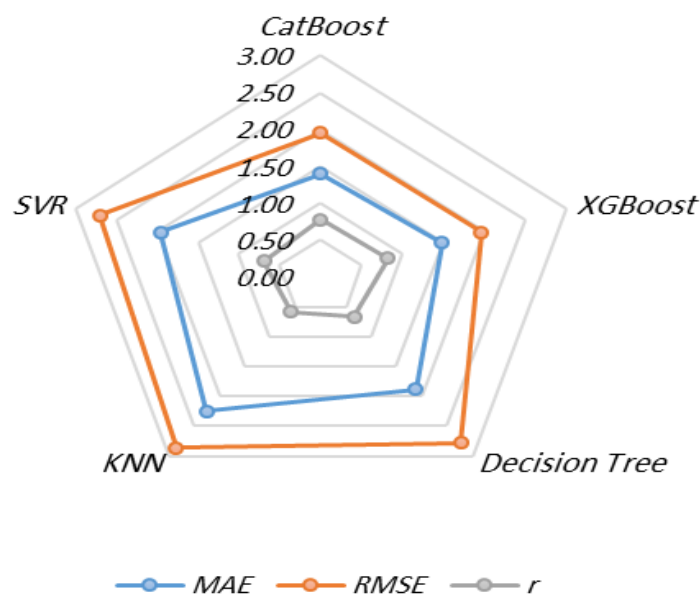
درصد بهبود عملکرد در شاخص‌های خطا داشته‌اند. این میزان بهبود حاکی از توانایی قابل توجه این مدل‌ها در مدل‌سازی روابط چندمتغیره‌ی غیرخطی میان متغیرهای اقلیمی و انسانی و بازتولید نوسانات سطح آب زیرزمینی در دشت شاهرود و بسطام است.

شکل ۳ رابطه بین مقادیر مشاهده‌شده و پیش‌بینی‌شده سطح آب زیرزمینی بر روی داده تست برای مدل‌های مختلف از جمله CatBoost، XGBoost، ماشین بردار پشتیبان (SVR)، درخت تصمیم (DT) و (KNN) نشان می‌دهد. در این نمودار، هر نقطه بیانگر یک مقدار پیش‌بینی‌شده در برابر مقدار واقعی متناظر آن است و خط قرمز نشان‌دهنده خط برابری ($y = x$) می‌باشد که

جدول ۴. عملکرد مدل‌های یادگیری ماشین.

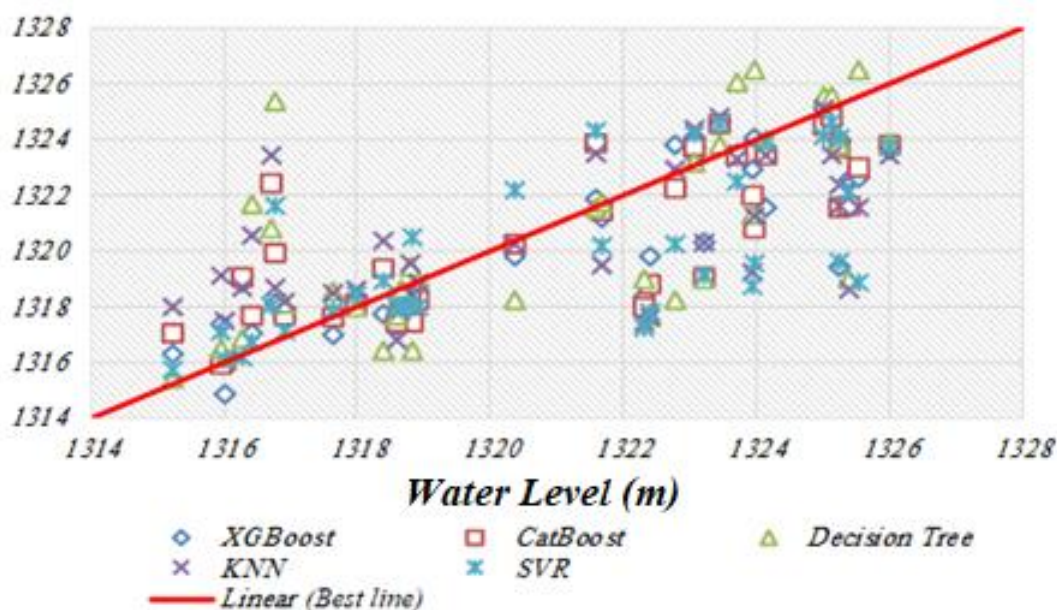
Table 4. Performance of Machine Learning Models.

مدل Model	MAE (m)	RMSE (m)	r
CatBoost	1.4029	1.9484	0.7815
XGBoost	1.4897	1.9692	0.8185
Decision Tree	1.8787	2.7790	0.6701
KNN	2.25	2.8617	0.5798
SVR	1.9584	2.6995	0.6903



شکل ۴. نمودار شاخص‌های ارزیابی مدل‌های یادگیری ماشین مورد استفاده در پژوهش حاضر.

Fig 4. Chart of Evaluation Metrics For The Machine Learning Models Used in The Present Study.



شکل ۵. نمودار پراکنش بین مقادیر پیش‌بینی شده و مشاهده شده سطح آب زیرزمینی توسط مدل‌های مختلف یادگیری ماشین.

Fig 5. Scatter Plot of Predicted Versus Observed Groundwater Levels By Different Machine Learning Models.

مدل برای پیش‌بینی تغییرات سطح آب زیرزمینی در این دشت است. این مدل می‌تواند به‌عنوان ابزاری کارآمد برای پایش و مدیریت منابع آب زیرزمینی در مناطق خشک و نیمه‌خشک مورد استفاده قرار گیرد. با به‌کارگیری داده‌های بیشتر از سال‌های اخیر و ترکیب این مدل‌ها با شبکه‌های عصبی یا مدل‌های هیبریدی، می‌توان دقت پیش‌بینی را بیشتر بهبود داد.

نتایج پژوهش حاضر در پیش‌بینی سطح آب زیرزمینی دشت شاهرود با یافته‌های Ahmadi و همکاران (۲۰۲۲) و Feng و همکاران (۲۰۲۴) هم‌خوانی قابل‌توجهی دارد، هرچند تفاوت‌هایی در مقیاس خطا و دقت مدل‌ها مشاهده می‌شود. در این مطالعه، مدل‌های XGBoost و CatBoost به‌ترتیب با ضریب همبستگی 0.818 و 0.7815 و مقادیر ریشه میانگین مربعات خطا (RMSE) به‌ترتیب $1/9692$ متر و $1/9484$ متر بهترین عملکرد را داشتند که برتری روش‌های Boosting در شناسایی روابط غیرخطی میان پارامترهای اقلیمی و انسانی را تأیید می‌کند. مطالعه Ahmadi و همکاران (۲۰۲۲) نشان داد دقت مدل‌های

در منطقه مورد مطالعه (دشت شاهرود و بسطام)، برداشت بیش از حد از منابع آب زیرزمینی طی سال‌های ۲۰۰۰ تا ۲۰۱۴ موجب افت تدریجی سطح آب شده است. بنابراین مدل‌هایی که بتوانند تأثیر توأمان عوامل طبیعی (مانند بارش و دما) و انسانی (مانند برداشت از چاه‌ها) را در نظر بگیرند، دقت بالاتری در پیش‌بینی خواهند داشت. مقدار بالای $r = 0.818$ در مدل XGBoost نشان‌دهنده توانایی آن در بازنمایی الگوی کلی تغییرات سطح آب در طول زمان است. در مقابل، مقادیر RMSE کمتر از ۲ متر در هر دو مدل CatBoost و XGBoost نشان می‌دهد که این مدل‌ها حتی در مقیاس مطلق نیز خطای قابل‌قبولی دارند. مدل‌های کلاسیک‌تر مانند KNN و DT، که بر اساس قوانین محلی یا تقسیم‌بندی ساده داده‌ها عمل می‌کنند، در این نوع داده‌های پیچیده عملکرد مطلوبی ندارند. مدل KNN به‌ویژه به دلیل وابستگی زیاد به ساختار داده و حساسیت به نویز، نتوانسته است روندهای بلندمدت را بازتاب دهد. در مجموع، می‌توان نتیجه گرفت که XGBoost با دقت و همبستگی بالاتر، مناسب‌ترین

یادگیری ماشین بیشتر به کیفیت داده‌ها، طول بازه زمانی و انتخاب متغیرهای کلیدی مانند بارش و دما بستگی دارد تا نوع الگوریتم، که با داده‌های ۱۵ ساله‌ی پژوهش حاضر هم‌راستا است. در مقابل، Feng و همکاران (۲۰۲۴) با مدل‌های یادگیری عمیق مانند CNN و RNN به دقت بالاتری (ضریب تعیین ۰/۹۹۵ و خطای RMSE برابر ۰/۰۵۶) دست یافتند، اما تفاوت در تفکیک‌پذیری زمانی (روزانه) و نرمال‌سازی داده‌ها علت اصلی این اختلاف است. بنابراین، اختلاف ظاهری در دقت مدل‌ها بیشتر ناشی از تفاوت در مقیاس و شیوه‌ی پردازش داده‌هاست تا برتری الگوریتمی. به‌طور کلی، نتایج این پژوهش همانند Ahmadi بر اهمیت داده‌های باکیفیت و بازه‌ی زمانی کافی تأکید دارد و مشابه Feng نشان می‌دهد مدل‌های پیشرفته مانند Boosting یا شبکه‌های عمیق می‌توانند دقت پیش‌بینی سطح آب زیرزمینی را بهبود دهند. در مجموع، مدل‌های XGBoost و CatBoost با عملکردی پایدار در داده‌های ماهانه، گزینه‌ای کارآمد و کم‌هزینه برای مدل‌سازی منابع آب زیرزمینی به‌شمار می‌آیند، در حالی که مدل‌های هیبریدی آینده می‌توانند به دقت شبکه‌های عمیق نزدیک شوند.

در برخی مطالعات، مدل‌هایی مانند SVR یا شبکه‌های عصبی بهترین عملکرد را نشان داده‌اند؛ با این حال در دشت شاهرود، مدل‌های Boosting شامل CatBoost و XGBoost نتایج دقیق‌تری ارائه دادند. علت اصلی این تفاوت به ماهیت غیرخطی، ناپایستایی و پیچیده داده‌های منطقه بازمی‌گردد. ترکیب عوامل اقلیمی و انسانی، نوسانات شدید برداشت، و ناهمگنی هیدروژئولوژیکی شاهرود سبب شده است که روابط میان متغیرها پیچیده‌تر از آن باشد که مدل‌های کلاسیک مانند SVR، KNN یا درخت تصمیم بتوانند آن را به‌طور کامل بازنمایی کنند. از سوی دیگر، الگوریتم‌های Boosting با بهره‌گیری از ساختار تجمیعی درخت‌های متعدد و توانایی بالا در

استخراج الگوهای غیرخطی، سازگاری بیشتری با چنین داده‌هایی دارند. مدل‌هایی مانند SVR برای عملکرد مطلوب نیازمند تنظیم دقیق پارامترها هستند و در حضور نویز و ناهمگنی شدید عملکرد آن‌ها افت می‌کند. بنابراین، برتری Boosting در این مطالعه ناشی از توان بالای آن در مدل‌سازی وابستگی‌های پیچیده و متغیرهای اثرگذار متعدد در دشت شاهرود است.

نتیجه‌گیری

با توجه به نقش حیاتی منابع آب زیرزمینی در تأمین نیازهای شرب، کشاورزی و صنعتی، و همچنین پیامدهای گسترده کاهش این منابع بر اکوسیستم‌ها و توسعه پایدار، بررسی و مدل‌سازی تغییرات سطح آب زیرزمینی از اهمیت ویژه‌ای برخوردار است. در پژوهش حاضر با هدف تحلیل تغییرات سطح آب زیرزمینی و ارزیابی عملکرد پنج الگوریتم یادگیری ماشین شامل CatBoost، XGBoost، درخت تصمیم (DT)، ماشین بردار پشتیبان (SVR) و K-نزدیک‌ترین همسایه (KNN) در دشت شاهرود و بسطام، از داده‌های اقلیمی و انسانی نظیر بارش، دما، میزان برداشت از چاه‌ها و قنات‌ها، و آب برگشتی کشاورزی در بازه زمانی ۲۰۰۰ تا ۲۰۱۴ استفاده شد. داده‌ها با نسبت ۸۰ درصد برای آموزش و ۲۰ درصد برای آزمون مورد استفاده قرار گرفتند.

نتایج ارزیابی بر اساس شاخص‌های MAE، RMSE و ضریب همبستگی (r) نشان داد که مدل‌های مبتنی بر تقویت گرادیان (Boosting)، شامل CatBoost و XGBoost، عملکرد برتری نسبت به سایر مدل‌ها دارند. در مدل CatBoost، مقدار MAE برابر با ۱/۴۰۲۹ و RMSE برابر با ۱/۹۴۸۴ به دست آمد که کمترین میزان خطا را در میان مدل‌های بررسی‌شده نشان می‌دهد. همچنین، در مدل XGBoost مقدار ضریب همبستگی (r) برابر با ۰/۸۱۸۵ محاسبه شد که بیانگر بالاترین میزان تطابق

F. (2016). Artificial neural network architectures and training processes. In *Artificial Neural Networks: A Practical Course* (pp. 21–28). Springer.

<https://doi.org/10.1016/j.jhydrol.2013.06.037>

Dastourani, M. T., Khayat, A., & Akhondi, Z. (2025). Evaluation of Artificial Intelligence Models' Accuracy for Groundwater Level Prediction (Case Study: Jahrom Plain). *Journal of Aquifer and Qanat*, 19(1), 67–82.

<https://doi.org/10.22077/jaaq.2025.9046.1105>

Ebrahimi, A., Pourdeilami, A., Khoshnevisan, S., & Asli Charandabi, M. (2025). Comparative Evaluation of Machine Learning Algorithms for Evaporation Estimation in Shahrood Region. *Journal of Hydraulic and Water Engineering*. 2025. <https://doi.org/10.224/JHWE.2025.16384.1070>.

Eftekhari, M., Haji Elyasi, A., & Eslami Nejad, S. A. (2024). Evaluation of Machine Learning Models in GIS for Groundwater Prediction in Arid Regions of Eastern Iran (Birjand Plain). *Journal of Aquifer and Qanat*, 18(2), 145–160.

<https://doi.org/10.22077/jaaq.2024.7282.1062>

El Ansari, R., El Bouhadioui, M., Aboutafail, M. O., Mejjad, N., Jamil, H., Jamal, E., Rissouni, Y., Zouiten, M., Boutracheh, H., & Moumen, A. (2023). A review of Machine learning models and parameters for groundwater issues. *Proceedings of the 6th International Conference on Networking, Intelligent Systems & Security*, <https://doi.org/10.1145/3607720.3607777>

Emamgholizadeh, S., Kashi, H., Marofpoor, I., & Zalaghi, E. (2014a). Prediction of water quality parameters of Karoon River (Iran) by artificial intelligence-based models. *International Journal of Environmental Science and Technology*, 11(3), 645–656. <https://doi.org/10.1007/s13762-013-0231-9>

Emamgholizadeh, S., Moslemi, K., & Karami, G. (2014b). Predicting the groundwater level of Bastam plain (Iran) by artificial neural network (ANN) and adaptive neuro-fuzzy inference system (ANFIS). *Water resources management*, 28(15), 5433–5446. <https://doi.org/10.1007/s11269-014-0814-1>

Feng, F., Ghorbani, H., & Radwan, A. E. (2024). Predicting groundwater level using traditional and deep machine learning algorithms. *Frontiers in Environmental Science*, 12, 1291327. <https://doi.org/10.3389/fenvs.2024.1291327>

Hsu, K. L., Gupta, H. V., & Sorooshian, S. (1995). Artificial neural network modeling of the rainfall-runoff process. *Water Resources Research*, 31(10), 2517–2530. <https://doi.org/10.1029/95WR01955>

Hussein, E. A., Thron, C., Ghaziasgar, M., Bagula, A., & Vaccari, M. (2020). Groundwater prediction using machine-learning tools. *Algorithms*, 13(11), 300. <https://doi.org/10.3390/a13110300>

میان مقادیر واقعی و پیش‌بینی‌شده و توانایی مناسب این مدل در بازنمایی روند زمانی تغییرات سطح آب زیرزمینی است. مقایسه کمی شاخص‌ها نشان می‌دهد که مدل‌های Boosting به طور میانگین بین ۲۵ تا ۳۵ درصد کاهش خطا و حدود ۲۰ تا ۲۵ درصد افزایش ضریب همبستگی نسبت به مدل‌های کلاسیک‌تر داشته‌اند. این برتری به توانایی این مدل‌ها در استخراج روابط غیرخطی میان متغیرهایی مانند بارش، دما و برداشت انسانی مرتبط است. در مقابل، مدل‌های ساده‌تری مانند KNN و درخت تصمیم به دلیل محدودیت در شناسایی الگوهای پیچیده، دقت پایین‌تری از خود نشان دادند.

به طور کلی، یافته‌های پژوهش بیانگر قابلیت بالای دو مدل XGBoost و CatBoost در مدل‌سازی رفتار سطح آب زیرزمینی است. با توجه به پتانسیل این مدل‌ها، پیشنهاد می‌شود در مطالعات آینده از داده‌های به‌روزتر و متغیرهای هیدروژئولوژیکی تکمیلی همچون کاربری اراضی و کیفیت آب استفاده شود. همچنین بهره‌گیری از رویکردهای ترکیبی مبتنی بر شبکه‌های عصبی عمیق و روش‌های Boosting می‌تواند دقت مدل‌سازی و توانایی شناسایی الگوهای پیچیده را افزایش دهد. در مجموع، به‌کارگیری الگوریتم‌های یادگیری ماشین می‌تواند زمینه‌ساز مدیریت هوشمندانه و پایدار منابع آب زیرزمینی و کاهش اثرات کمبود آب باشد.

منابع

Ahmadi, A., Olyaei, M., Heydari, Z., Emami, M., Zeynolabedin, A., Ghomlaghi, A., Daccache, A., Fogg, G. E., & Sadegh, M. (2022). Groundwater level modeling with machine learning: a systematic review and meta-analysis. *Water*, 14(6), 949. <https://doi.org/10.1038/s41467-023-42411-2>

Amanambu, A. C., Obarein, O. A., Mossa, J., Li, L., Ayeni, S. S., Balogun, O., Oyebamiji, A., & Ochege, F. U. (2020). Groundwater system and climate change: Present status and future considerations. *Journal of Hydrology*, 589, 125163. <https://doi.org/10.1016/j.jhydrol.2020.125163>

Da Silva, I. N., Hernane Spatti, D., Andrade Flauzino, R., Liboni, L. H. B., & dos Reis Alves, S.

- Klove, B., Ala-Aho, P., Bertrand, G., Gurdak, J. J., Kupfersberger, H., Kværner, J., Muotka, T., Mykrä, H., Preda, E., & Rossi, P. (2014). Climate Change Impacts on Groundwater and Dependent Ecosystems-in press.
<https://doi.org/10.1016/j.jhydrol.2013.06.037>
- Nohani, E., Babaali, H., & Dehghani, R. (2025). Evaluation of Meta-Heuristic Models' Accuracy in Groundwater Level Analysis of Delfan Plain. *Journal of Aquifer and Qanat*, 19(2), 101–118.
<https://doi.org/10.22077/jaaq.2025.8834.1096>
- Noori, R., Maghrebi, M., Jessen, S., Bateni, S., Heggy, E., Javadi, S., Noury, M., Pistre, S., Abolfathi, S., & AghaKouchak, A. (2023). Decline in Iran's groundwater recharge, *Nat. Commun.*, 14, 6674. <https://doi.org/10.1038/s41467-023-42411-2>
- Rahman, A. S., Hosono, T., Quilty, J. M., Das, J., & Basak, A. (2020). Multiscale groundwater level forecasting: Coupling new machine learning approaches with wavelet transforms. *Advances in Water Resources*, 141, 103595.
<https://doi.org/10.1016/j.advwatres.2020.103595>
- Rezapour, A., & Sabzekar, M. (2025). An Intelligent Combination of Machine Learning Approaches for Groundwater Fluctuations Prediction. *Iranian Journal of Science and Technology, Transactions of Civil Engineering*, 1–15. <https://doi.org/10.1007/s40996-025-01742-4>
- Singh, A., Patel, S., Bhadani, V., Kumar, V., & Gaurav, K. (2024). AutoML-GWL: an automated machine learning model for the prediction of groundwater level. *Engineering Applications of Artificial Intelligence*, 127, 107405.
<https://doi.org/10.1016/j.engappai.2024.107405>
- Zhao, W. G., Wang, H., & Wang, Z. J. (2011). Groundwater level forecasting based on a support vector machine. *Applied Mechanics and Materials*, 44, 1365–1369.
<https://doi.org/10.4028/www.scientific.net/AMM.44.1365>
- Zhu, F., Sun, Y., Han, M., Hou, T., Zeng, Y., Lin, M., Wang, Y., & Zhong, P.-a. (2025). A robust Bayesian multi-machine learning ensemble framework for probabilistic groundwater level forecasting. *Journal of Hydrology*, 650, 132567.
<https://doi.org/10.1016/j.jhydrol.2025.132567>