

Monthly Rainfall Simulation Using Hybrid Machine Learning Models: NAMAK Lake Basin

A. DalirGabrabad¹, M.A. Abdollahi¹, A. Ashrafzadeh^{1*}

1- Department of Irrigation and Reclamation Engineering, College of Agriculture and Natural Resources, University of Tehran, Karaj, Iran

(*- Corresponding Author Email: a.ashrafzadeh@ut.ac.ir)

Received: 29 January 2026

Revised: 28 May 2026

Accepted: 01 June 2026

Available Online: 01 June 2026

How to cite this article:

DalirGabrabad, A., Abdollahi, M.A., & Ashrafzadeh, A. (2026). Monthly rainfall simulation using hybrid machine learning models: NAMAK lake basin. *Journal of Water and Soil*, 39(6), 553-572. (In Persian with English abstract). <https://doi.org/10.22067/jsw.2026.97595.1522>

Introduction

Precipitation is a fundamental component of the hydrological cycle and plays a vital role in water resource management, agriculture, ecosystem sustainability, hydropower generation, disaster mitigation, and urban planning. Accurate rainfall prediction is essential for effective water allocation, flood and drought risk management, and the development of early warning systems, thereby reducing potential damages to human communities and infrastructure. In agriculture, reliable precipitation forecasts support optimal planting schedules and resource management, leading to increased productivity and reduced weather-related losses. Furthermore, rainfall prediction is crucial for reservoir operation and sustainable long-term water resource planning.

Due to the limited spatial coverage of ground-based meteorological stations, satellite-based precipitation products have gained increasing attention as reliable alternatives for regional-scale rainfall analysis. Among these, the Climate Hazards Group InfraRed Precipitation with Stations (CHIRPS) dataset is widely recognized for its extensive spatial coverage, relatively high accuracy in arid and mountainous regions, and free accessibility. CHIRPS integrates satellite observations with ground station data and provides precipitation estimates at approximately 5 km spatial resolution across multiple temporal scales since 1981. In this study, monthly mean CHIRPS precipitation data for the period 2000–2024 were extracted and analyzed for rainfall prediction purposes.

The inherent nonlinearity and complex behavior of precipitation processes make traditional statistical approaches insufficient for accurate forecasting. Consequently, machine learning techniques have emerged as powerful tools for modeling complex climatic phenomena. Although numerous previous studies have applied machine learning algorithms, such as Artificial Neural Networks (ANN), Random Forest (RF), Support Vector Machines (SVM), and Convolutional Neural Networks (CNN), to rainfall prediction, most have focused on single-model evaluations or simple model comparisons. However, combining multiple models has been shown to reduce prediction uncertainty and improve forecast robustness, an approach that has received comparatively limited attention.

The primary objective of this research is to evaluate the performance of satellite-based CHIRPS precipitation data and to enhance monthly rainfall prediction using machine learning models implemented in the WEKA environment. To achieve this goal, fifteen individual machine learning models were first assessed independently. Subsequently, pairwise combinations of these models (105 combinations in total) were developed using simple averaging and brute-force weighting strategies to improve predictive accuracy. The novelty of this study lies not merely in the application of multiple machine learning algorithms but in proposing a systematic, transparent, and reproducible framework for satellite-based rainfall prediction in data-scarce basins. This framework is region-independent and can be readily applied to other climatic and hydrological settings.



Authors retain the copyright. This is an open access article distributed under [Creative Commons Attribution 4.0 International License \(CC BY 4.0\)](https://creativecommons.org/licenses/by/4.0/).

<https://doi.org/10.22067/jsw.2026.97595.1522>

Materials and Methods

This study investigates monthly precipitation prediction in the Salt Lake Basin using satellite-based CHIRPS data for the period January 2000 to December 2025, extracted via the Google Earth Engine platform. CHIRPS was selected due to its adequate spatial resolution (0.05°), integration of satellite and ground-based observations, long-term temporal coverage since 1981, and suitability for arid and semi-arid regions with sparse rain-gauge networks. Data quality was assessed using the Interquartile Range (IQR) method, and no outliers were detected. As the dataset exhibited consistent scaling across all time steps, no normalization was applied. The data were divided into training (January 1999–December 2021) and testing (January 2022–December 2024) subsets. All modeling procedures were implemented in the WEKA environment using time-series modules. To validate the CHIRPS dataset, satellite-derived precipitation estimates were compared with observations from four ground stations over ten years (2007–2017), and basin-scale performance was evaluated using averaged statistical metrics. Initially, sixteen machine learning models were examined, and one model (Random Tree) was excluded due to poor performance. The final set of fifteen models included Gaussian Processes, Random Forest, MLP Regressor, RBF Regressor, Linear Regression, SMOreg, IBk, LWL, Additive Regression, Bagging, Random Committee, Random SubSpace, Decision Table, M5Rules, and M5P. Model hyperparameters were optimized using the Grid Search tool in WEKA. To improve prediction accuracy and robustness, pairwise ensemble models were generated using simple averaging, resulting in 105 combined configurations. Additionally, a brute-force weighting approach was applied by assigning weights between 0 and 1 with a step of 0.1 to each model, subject to a unity-sum constraint, in order to identify the optimal ensemble structure.

Model performance was evaluated using the correlation coefficient, RMSE, MAE, MSE, bias, and Nash–Sutcliffe efficiency (NSE), and the most accurate configuration was selected based on testing results.

Results and Discussion

The accuracy of the CHIRPS satellite precipitation product was first evaluated using observed rainfall data from four rain-gauge stations distributed across the Salt Lake Basin. The validation results indicated a satisfactory agreement between satellite-derived and observed precipitation, with an overall coefficient of determination (R^2) of approximately 0.69 and an NSE of 0.70, confirming the suitability of CHIRPS data for monthly rainfall analysis in the study area. Subsequently, the performance of fifteen individual machine learning models was assessed for monthly precipitation prediction using the testing period (January 2022–December 2024). Model evaluation based on correlation coefficient, RMSE, MAE, MSE, bias, and Nash–Sutcliffe efficiency (NSE) revealed notable differences among algorithms. Rule-based and tree-based models, particularly M5Rules and Additive Regression, exhibited superior performance, characterized by higher NSE values, lower error magnitudes, and relatively small bias. In contrast, instance-based models such as IBk and function-based models like RBF Regressor showed weaker performance, with larger error dispersion and lower correlation with observed data. The combined analysis of bias, residual standard deviation, and median absolute error provided a more comprehensive understanding of model behavior than single error metrics alone. Visual assessments using Taylor diagrams and scatter plots further confirmed the robustness of M5Rules, followed by Additive Regression and Random Forest, in capturing both the variability and temporal patterns of observed precipitation.

To enhance prediction accuracy and stability, pairwise combinations of the individual models were developed using simple averaging, resulting in 105 ensemble configurations. The results demonstrated that most ensemble models outperformed their corresponding single-model counterparts, as evidenced by higher NSE values and reduced error metrics. The combination of Additive Regression–M5Rules achieved the best overall performance, yielding the lowest MAE and MSE and the highest NSE among all tested configurations. Histogram analysis of NSE values showed that a large proportion of the ensemble models achieved NSE values above 0.75, indicating the general effectiveness of the ensemble approach rather than improvement limited to specific combinations. In addition, a brute-force weighting strategy was applied to optimize model contributions within the ensembles. Compared to simple averaging, the Brute Force approach consistently improved NSE values by reducing the influence of weaker models and assigning higher weights to more accurate ones. This improvement was particularly evident in combinations involving low-performing models such as IBk and RBF Regressor. Nevertheless, simple averaging also produced competitive and stable results without increasing computational complexity.

Overall, the results indicate that monthly rainfall prediction accuracy depends not only on the choice of individual algorithms but also on the interaction and complementarity of model errors. The proposed ensemble framework effectively reduced bias and variance while avoiding overfitting through strict temporal data separation and independent testing. These findings demonstrate that systematic pairwise model combination, even with simple averaging, can substantially enhance the accuracy and robustness of precipitation forecasts in data-scarce and climatically complex regions.

Conclusion

The Salt Lake Basin is characterized by sparse rain-gauge coverage and high climatic variability, making accurate monthly precipitation estimation essential for water resources management. This study evaluated the CHIRPS satellite precipitation dataset and assessed the performance of individual and combined machine learning models for monthly rainfall prediction. Validation against four ground stations confirmed that CHIRPS reliably represents the temporal pattern of monthly precipitation in the basin, with acceptable correlation and Nash–Sutcliffe efficiency values, indicating its suitability for hydrological applications in arid and semi-arid regions. Among the fifteen evaluated machine learning models, M5Rules, Additive Regression, and Random Forest exhibited superior performance as single models. Ensemble modeling further improved prediction accuracy by reducing systematic errors and enhancing robustness. The Additive Regression–M5Rules combination achieved the best performance ($NSE \approx 0.88$), outperforming all individual models. These results highlight the effectiveness of integrating CHIRPS data with ensemble machine learning approaches for improving monthly precipitation prediction in data-scarce basins. The proposed framework provides a reliable and transferable tool for supporting sustainable water resources management in arid and semi-arid regions.

Keywords: Climate data prediction, Precipitation data modeling, WEKA

شبیه‌سازی بارش ماهانه با مدل‌های ترکیبی یادگیری ماشین: حوضه دریاچه نمک

علی دلیر گبرآباد^۱ - محمدامین عبداللهی^۱ - افشین اشرف‌زاده^{۱*}

تاریخ دریافت: ۱۴۰۴/۱۱/۰۹

تاریخ پذیرش: ۱۴۰۵/۰۳/۱۱

چکیده

پیش‌بینی دقیق بارش ماهانه به‌ویژه در مناطق خشک و نیمه‌خشک، برای مدیریت منابع آب و برنامه‌ریزی‌های هیدرولوژیک اهمیت بالایی دارد. با توجه به کمبود و پراکندگی ایستگاه‌های زمینی در حوضه دریاچه نمک، استفاده از داده‌های ماهواره‌ای با دقت مکانی مناسب، ضرورت می‌یابد. نوآوری اصلی این پژوهش، ارزیابی کارایی داده‌های بارش CHIRPS در مدل‌سازی بارش ماهانه با بهره‌گیری از طیفی متنوع از مدل‌های یادگیری ماشین و همچنین بررسی اثربخشی مدل‌های ترکیبی در افزایش دقت پیش‌بینی است. در این مطالعه، داده‌های بارش CHIRPS با تفکیک افقی 0.5° برای دوره ۲۰۰۰ تا ۲۰۲۴ استخراج و پس از اعتبارسنجی اولیه مورد استفاده قرار گرفت. در این پژوهش برای ارزیابی داده‌های بارش CHIRPS، طیفی از مدل‌های یادگیری ماشین شامل مدل‌های درختی، مدل‌های رگرسیونی، مدل‌های کرنلی، مدل‌های مبتنی بر فاصله و مدل‌های ترکیبی به‌منظور مدل‌سازی رفتار غیرخطی بارش ماهانه مورد استفاده قرار گرفتند. این تنوع رویکردها، امکان ارزیابی تطبیقی ساختارهای مختلف یادگیری و بررسی تأثیر معماری مدل بر دقت پیش‌بینی را فراهم می‌سازد. مدل M5Rules با بدست آوردن $NSE = 0.80$ بهترین عملکرد را در مدل‌های منفرد داشت. همچنین برای بهبود دقت، بر پایه میانگین‌گیری ساده، 105 مدل ترکیبی به‌صورت دوبه‌دو تولید و ارزیابی شد و سپس برای پیدا کردن بهترین وزن از روش brute force (جستجوی فراگیر برای وزن‌دهی بهینه) استفاده شد. نتایج نشان داد که؛ مدل‌های ترکیبی به‌طور کلی عملکرد بهتری نسبت به مدل‌های منفرد دارند. ترکیب Additive Regression-M5Rules بهترین عملکرد را نسبت به سایر ترکیب‌ها با $MAE = 0.155$ ، $MSE = 0.05$ و $NSE = 0.88$ ارائه کرد. این یافته‌ها نشان می‌دهد که ترکیب داده‌های CHIRPS با مدل‌های ترکیبی یادگیری ماشین می‌تواند رویکردی کارآمد و قابل اتکا برای پیش‌بینی بارش ماهانه در مناطق با شبکه ایستگاهی محدود باشد.

واژه‌های کلیدی: الگوریتم‌های یادگیری ماشین، پیش‌بینی داده‌های اقلیمی، مدل‌سازی داده‌های بارش، وکا

مقدمه

برف به باران و افزایش مشکلات مربوط به سیل و آبگرفتگی به ویژه در مناطق شهری، پیش‌بینی‌های قابل اعتماد امکان ایجاد سیستم‌های هشدار زودهنگام را فراهم می‌کند که برای اقدامات پیشگیرانه جهت حفاظت از جوامع و دارایی‌ها ضروری هستند (Abdollahi et al., 2024; Fikileni & Wolski, 2022). پیش‌بینی بارش در کشاورزی، به‌ویژه در کشورهایی مانند هند، برای تعیین زمان مناسب کشت و مدیریت منابعی چون آب و کود ضروری است. این اقدامات باعث افزایش بهره‌وری، بهبود عملکرد محصولات، و کاهش خسارات ناشی از شرایط جوی نامساعد می‌شود (Patel et al., 2023). پیش‌بینی بارش در مدیریت مخازن و سدها نقش کلیدی دارد و با تنظیم ذخیره و آزادسازی آب، نیازهای کشاورزی، صنعتی و خانگی را برآورده می‌سازد. همچنین، داده‌های پیش‌بینی کمک می‌کنند تا منابع آب به‌صورت پایدار، مدیریت شده و دسترسی طولانی‌مدت به آن تضمین

بارش یکی از مهم‌ترین متغیرهای چرخه هیدرولوژیکی است که؛ نقش حیاتی در مدیریت منابع آب، کشاورزی و اکوسیستم‌های طبیعی ایفا می‌کند. پیش‌بینی دقیق میزان بارش نه تنها برای برنامه‌ریزی منابع آب ضروری است؛ بلکه نقش مهمی در مدیریت خطرات سیل و خشکسالی دارد. زیرا کاهش آسیب‌های احتمالی به انسان‌ها و زیرساخت‌ها را در پی دارد. همچنین با توجه به تغییر نوع بارش‌ها از

۱- گروه مهندسی آبیاری و آبادانی، دانشکده‌گان کشاورزی و منابع طبیعی دانشگاه تهران، کرج، ایران

(*)- نویسنده مسئول: (Email: a.ashrafzadeh@ut.ac.ir)

<https://doi.org/10.22067/jsw.2026.97595.1522>

علمی و کاربردی مشخصی صورت گرفته است. نخست، CHIRPS دارای پوشش زمانی بلندمدت و پیوسته از سال ۱۹۸۱ میلادی تاکنون است که امکان تحلیل روندهای اقلیمی در مقیاس چنددهه‌ای را فراهم می‌کند؛ در حالی که برخی محصولات ماهواره‌ای مانند TRMM و GPM دوره زمانی کوتاه‌تری را پوشش می‌دهند (Funk et al., 2015). دوم، این مجموعه داده با ترکیب برآوردهای ماهواره‌ای و داده‌های ایستگاه‌های زمینی، دقت خود را به‌ویژه در مناطق خشک و کوهستانی افزایش داده است؛ که این ویژگی برای منطقه مورد مطالعه اهمیت بالایی دارد (Beck et al., Funk et al., 2015). علاوه بر این، وضوح مکانی مناسب (حدود ۵ کیلومتر) و ارائه داده‌ها در مقیاس‌های زمانی روزانه، دهه‌ای و ماهانه، انعطاف‌پذیری لازم برای تحلیل‌های آماری و مکانی را فراهم می‌سازد. از نظر اجرایی نیز، دسترسی رایگان، به‌روزرسانی منظم، ارائه در فرمت‌های استاندارد (مانند NetCDF^{۱۱} و GeoTIFF^{۱۲}) و سازگاری کامل با نرم‌افزارهای متداول تحلیل داده، استفاده از این مجموعه داده را تسهیل می‌کند و نیاز به پیش‌پردازش‌های پیچیده را کاهش می‌دهد. همچنین، کاربرد گسترده CHIRPS در مطالعات منطقه‌ای و بین‌المللی موجب افزایش قابلیت مقایسه‌پذیری نتایج این پژوهش با سایر تحقیقات شده است. بر این اساس، انتخاب CHIRPS نه تنها بر مبنای دسترسی، بلکه بر اساس ترکیب پوشش زمانی بلندمدت، دقت مکانی مناسب، ادغام داده‌های زمینی و ماهواره‌ای و اعتبار علمی گسترده آن صورت گرفته است. در این تحقیق، داده‌های میانگین ماهانه CHIRPS برای بازه زمانی ۲۰۰۰ تا ۲۰۲۴ میلادی به‌منظور بررسی پیش‌بینی بارش در منطقه مورد مطالعه استخراج و تحلیل شده‌اند.

با توجه به نوسانات اقلیمی و اثرات فزاینده آن بر منابع آب، نیاز به پیش‌بینی مناسب‌تر بارش در دوره‌های زمانی مختلف، بیش از پیش احساس می‌شود. در این میان، داده‌های ماهواره‌ای به دلیل پوشش مکانی وسیع و پیوستگی زمانی، ابزاری مناسب برای استفاده در پژوهش‌ها در کنار داده‌های ایستگاهی محسوب می‌شوند. همچنین، با رشد الگوریتم‌های یادگیری ماشین، امکان مدل‌سازی دقیق‌تری از رفتار پیچیده و غیرخطی بارش فراهم شده است. در بسیاری از مطالعات پیشین، از مدل‌های یادگیری ماشین برای پیش‌بینی بارش استفاده شده، اما اغلب آن‌ها به ارزیابی یک یا چند مدل به‌صورت مجزا محدود شده‌اند. این در حالی است که؛ ترکیب خروجی مدل‌ها می‌تواند منجر به کاهش خطا و افزایش پایداری پیش‌بینی‌ها شود (Tuysuzoglu et al., 2023)، موضوعی که تاکنون کمتر مورد توجه

شود (Ramani et al., 2024). پیش‌بینی بارش برای تولید برق آبی ضروری است، زیرا با پیش‌بینی جریان آب به مخازن، تولید برق را بهینه و عرضه پایدار انرژی را تضمین می‌کند. همچنین، این پیش‌بینی‌ها تخصیص بهینه منابع انرژی را ممکن کرده و به کاهش خطر کمبود برق کمک می‌کنند، پیش‌بینی بارش در برنامه‌ریزی شهری برای طراحی سیستم‌های زهکشی کارآمد و کاهش خطر آبگرفتگی شهری مؤثر است. همچنین، این پیش‌بینی‌ها در توسعه زیرساخت‌هایی مانند جاده‌ها و پل‌ها با ارائه اطلاعات درباره شرایط جوی و خطرات احتمالی نقش مهمی دارند (Patel et al., 2023). پیش‌بینی دقیق بارش برای مدیریت مؤثر منابع آب، برنامه‌ریزی کشاورزی، آمادگی در برابر بلایا و توسعه زیرساخت ضروری است (Salahi et al., 2024a,b). بهبود مستمر در فناوری‌های پیش‌بینی برای پاسخگویی به چالش‌های رو به رشد ناشی از تغییرات آب و هوایی و اطمینان از مدیریت پایدار منابع آب ضروری است.

با توجه به عدم کفایت پوشش مکانی ایستگاه‌های زمینی، داده‌های ماهواره‌ای مانند CHIRPS^۱ به‌عنوان منابع قابل اعتماد برای تحلیل بارش در مقیاس‌های منطقه‌ای مورد توجه قرار گرفته‌اند. از سوی دیگر، با پیچیدگی و ماهیت غیرخطی الگوهای بارش، مدل‌های یادگیری ماشین^۲ ابزارهای مناسبی برای تحلیل چنین داده‌هایی محسوب می‌شوند. داده‌های سنجش از دور با فناوری‌های پیشرفته، نقشی کلیدی در پایش محیط زیست، برنامه‌ریزی شهری و مدیریت منابع آب دارند. این داده‌ها، امکان پایش بلادرنگ تغییرات محیطی و جنگل‌زدایی را فراهم می‌سازند (Mahjoobi & Rafiei, 2025). تهیه نقشه‌ها و تحلیل تنوع زیستی (Fagan & DeFries, 2024)، رصد دقیق تغییرات کاربری اراضی با تصاویر وضوح بالا و ادغام با GIS^۳ و یادگیری ماشین را ارائه می‌دهند (Xu & Zheng, 2024). با توجه به تنوع مجموعه داده‌های بارش در دسترس نظیر TRMM^۴، GPM^۵، MERRA-2^{۱۰}، ERA5^۶، CRU^۸، GPCC^۷، PERSIANN^۹، انتخاب مجموعه داده CHIRPS در این پژوهش بر اساس معیارهای

- 1- CHIRPS: Climate Hazards Group InfraRed Precipitation with Station data
- 2- Machine Learning
- 3- Geographic Information System
- 4- Tropical Rainfall Measuring Mission
- 5- Global Precipitation Measurement
- 6- Precipitation Estimation from Remotely Sensed Information using Artificial Neural Networks
- 7- Global Precipitation Climatology Centre
- 8- Climatic Research Unit
- 9- ECMWF Reanalysis v5
- 10- Modern-Era Retrospective Analysis for Research and Applications, Version 2

11- Network Common Data Form

12- Geographic Tagged Image File Format

قرار گرفته است.

با پیشرفت الگوریتم‌های داده‌کاوی^۱ و یادگیری ماشین، استفاده از این رویکردها برای مدل‌سازی پدیده‌های اقلیمی، رشد قابل توجهی یافته است. با وجود گسترش مطالعات در زمینه پیش‌بینی بارش با استفاده از مدل‌های هوش مصنوعی، بیشتر پژوهش‌ها بر ارزیابی تک‌مدل‌ها یا مقایسه چند الگوریتم متمرکز بوده‌اند. در مقابل، کمتر مطالعه‌ای به‌صورت سیستماتیک به ترکیب مدل‌های^۲ یادگیری ماشین و ارزیابی هم‌زمان چند الگوریتم در چارچوبی یکپارچه و قابل تکرار پرداخته است. هدف اصلی این پژوهش، ارزیابی داده‌های بارش ماهواره‌ای CHRS و پیش‌بینی بارش ماهانه با استفاده از داده‌های ماهواره‌ای و مدل‌های یادگیری ماشین در محیط نرم‌افزار WEKA^۳ است. در این راستا، ابتدا عملکرد ۱۵ مدل یادگیری ماشین به‌صورت مجزا بررسی شده و سپس با ترکیب دو به دوی آن‌ها (در مجموع ۱۰۵ حالت)، تلاش شده است تا دقت پیش‌بینی‌ها به‌طور قابل توجهی بهبود یابد. نوآوری و ارزش افزوده این پژوهش صرفاً در به‌کارگیری مدل‌های یادگیری ماشین یا افزایش تعداد آن‌ها خلاصه نمی‌شود، بلکه در ارائه یک چارچوب سیستماتیک و قابل تکرار برای ارزیابی داده‌های CHIRPS و بهبود پیش‌بینی بارش ماهانه در حوضه‌های کم‌داده نهفته است.

در این مطالعه، کلیه مراحل شامل استخراج داده، پیش‌پردازش، تقسیم‌بندی زمانی، آموزش، آزمون و ترکیب مدل‌ها به‌صورت شفاف و قابل بازتولید تشریح شده‌اند. استفاده از ترکیب‌های دویدهو با میانگین‌گیری ساده و وزن‌دهی BruteForce یک رویکرد عمومی است؛ که می‌تواند بدون وابستگی به منطقه خاص، در سایر حوضه‌ها نیز پیاده‌سازی شود. از این‌رو، پژوهش حاضر نه‌تنها قابل تکرار است، بلکه می‌تواند به‌عنوان یک چارچوب مرجع برای مطالعات مشابه در سایر مناطق اقلیمی مورد استفاده قرار گیرد. مطالعات مرتبط با پیش‌بینی بارش و برآورد تغییرات آن را می‌توان در چند دسته اصلی طبقه‌بندی کرد. این دسته‌ها از روش‌های داده‌کاوی و یادگیری ماشین تا مدل‌های آماری ریزمقیاس‌نمایی و رویکردهای هیبریدی را دربرمی‌گیرند.

در نخستین گروه، پژوهش‌هایی قرار می‌گیرند که از روش‌های مبتنی بر قواعد و داده‌کاوی برای استخراج الگوهای تصمیم‌گیری استفاده کرده‌اند. برای مثال، Omary و همکاران در سال ۲۰۱۲ با معرفی الگوریتم DIAL^۴ توانستند قواعد انجمنی مؤثر در پیش‌بینی محلی آب‌وهوا را استخراج کنند (Omary et al., 2012). همچنین

کومار و اسواتی (2023Kumar & Swathi) با استفاده از الگوریتم J48 در محیط WEKA نشان دادند که درخت تصمیم^۵ قادر است قوانین پیش‌بینی بارش را با دقت قابل قبول استخراج کند.

گروه دوم شامل پژوهش‌هایی است که یادگیری عمیق و شبکه‌های عصبی را برای پیش‌بینی بارش به‌کار گرفته‌اند. بال و همکاران با بهره‌گیری از شبکه‌های عصبی کانولوشنی^۶ نشان دادند که این مدل‌ها در پیش‌بینی بارش روزانه عملکرد بهتری نسبت به روش‌های سنتی دارند (Abishek et al., 2017). خان و همکاران (۲۰۲۲) نیز با بررسی شبکه‌های عصبی و مدل‌های ترکیبی گزارش کردند که این روش‌ها می‌توانند دقت پیش‌بینی بارش را به‌طور چشمگیری افزایش دهند (Rahman et al., 2022). در همین چارچوب، برخی پژوهش‌ها به پیش‌بینی بارش فصلی با استفاده از سیستم استنباط فازی^۷ و شبکه عصبی مصنوعی پرداخته‌اند. در این مطالعات، بارش بازه دسامبر تا می در منطقه خراسان بزرگ بررسی شده و نشان داده شده است که؛ استفاده از ANN^۸ و مدل فازی می‌تواند در برآورد بارش فصلی مؤثر باشد. این دسته از پژوهش‌ها به‌طور مشخص به مقایسه‌ی دو رویکرد فازی و شبکه عصبی برای پیش‌بینی فصلی می‌پردازند (Fallah-Ghalhari et al., 2009). Babaei Hessar & Ghazavi, 2014

گروه سوم شامل پژوهش‌هایی است که؛ از الگوریتم‌های کلاسیک یادگیری ماشین استفاده کرده‌اند. سونی و همکاران (۲۰۱۸) با مقایسه^۹ KNN، SVM، درخت تصمیم و شبکه عصبی مصنوعی، نشان دادند که ANN عملکرد بهتری در پیش‌بینی بارش روزانه دارد (Amrehn et al., 2018). در پژوهشی دیگر، اوسوال (Oswal, 2019) با استفاده از داده‌های هواشناسی استرالیا و روش‌هایی نظیر رگرسیون لجستیک^{۱۰}، KNN^{۱۱}، درخت تصمیم و مدل‌های ترکیبی، بر نقش تعیین‌کننده انتخاب مدل مناسب و کیفیت پیش‌پردازش داده‌ها تأکید کرد. پاتل و همکاران (Patel et al., 2023) نیز با ارزیابی الگوریتم‌های مختلف داده‌کاوی در نرم‌افزار WEKA نشان دادند که Random Forest در پیش‌بینی بارش عملکرد بهتری ارائه می‌دهد. علاوه بر این، ولاسکو و همکاران (Velasco et al., 2022) با بهره‌گیری از مدل SVR^{۱۲} نشان دادند که انتخاب بهینه هسته و پارامترهای مدل نقش مهمی در افزایش دقت پیش‌بینی دارد.

5- Decision Tree

6- CNN

7- Fuzzy

8- Artificial Neural Network

9- Support Vector Machine

10- Logistic Regression

11- K-Nearest Neighbors

12- Support Vector Regression

1- Data Mining

2- ensemble

3- Waikato Environment for Knowledge Analysis

4- Differential Absorption Lidar

مواد و روش‌ها

داده‌های مورد استفاده

در این پژوهش، داده‌های بارش ماهانه CHIRPS مربوط به بازه زمانی ژانویه ۲۰۰۰ تا دسامبر ۲۰۲۵ برای حوضه آبریز دریاچه نمک مورد استفاده قرار گرفت (شکل ۱). مجموعه داده CHIRPS توسط گروه Climate Hazards Group توسعه یافته و با تفکیک مکانی ۰٫۰۵ درجه از سال ۱۹۸۱ تاکنون ارائه می‌شود. این داده‌ها از ترکیب تصاویر مادون قرمز ماهواره‌ای، اقلیم‌نگاشت‌های بلندمدت و داده‌های ایستگاه‌های باران‌سنجی تولید می‌شوند. در فرآیند تولید ابتدا بارش بر اساس دمای ابر در تصاویر ماهواره‌ای برآورد می‌شود، سپس با استفاده از داده‌های اقلیمی و ایستگاه‌های زمینی تصحیح و کالیبره می‌گردد و در ادامه از روش‌های آماری برای کاهش بایاس و بهبود دقت استفاده می‌شود. به دلیل ترکیبی بودن داده‌ها، پوشش مکانی مناسب و عملکرد قابل اعتماد در مناطق دارای شبکه ایستگاهی محدود، CHIRPS به یکی از معتبرترین منابع برآورد بارش در مناطق خشک و نیمه‌خشک تبدیل شده و گزینه‌ای مناسب برای تحلیل بارش در حوضه آبریز دریاچه نمک محسوب می‌شود (Funk et al., 2015).

پیش‌پردازش و تقسیم‌بندی داده‌ها

داده‌های استخراج‌شده بارش ماهانه پیش از مدل‌سازی تحت فرآیند پیش‌پردازش قرار گرفتند. به منظور شناسایی داده‌های پرت، از روش دامنه بین چارکی^۱ استفاده شد. در این روش، چارک اول (Q1) و چارک سوم (Q3) محاسبه شده و دامنه بین چارکی مطابق رابطه (۱) تعیین گردید. داده‌هایی که در دامنه (a,b) قرار نگیرند به عنوان داده پرت در نظر گرفته می‌شوند (Turkey, 1977).

$$IQR = Q3 - Q1 \quad (1)$$

$$(a) = Q1 - 1.5 * IQR \quad (2)$$

$$(b) = Q3 + 1.5 * IQR \quad (3)$$

نتایج بررسی نشان داد که هیچ مقدار خارج از این بازه در سری زمانی مشاهده نشد؛ بنابراین داده‌ها فاقد نقاط پرت آماری معنادار بودند. از آنجاکه روش IQR یک روش ناپارامتری است و مبتنی بر توزیع داده نمی‌باشد، آزمون سطح معناداری کلاسیک (مانند آزمون‌های پارامتری) در این مرحله موضوعیت نداشت و معیار شناسایی بر اساس قاعده $IQR \times 1/5$ انجام شد. همچنین به منظور بررسی ناسازگاری احتمالی در سری زمانی، روند کلی و پیوستگی مقادیر در بازه زمانی مورد مطالعه از نظر نوسانات غیرمنطقی یا گسست ناگهانی بررسی گردید؛ که موردی مشاهده نشد.

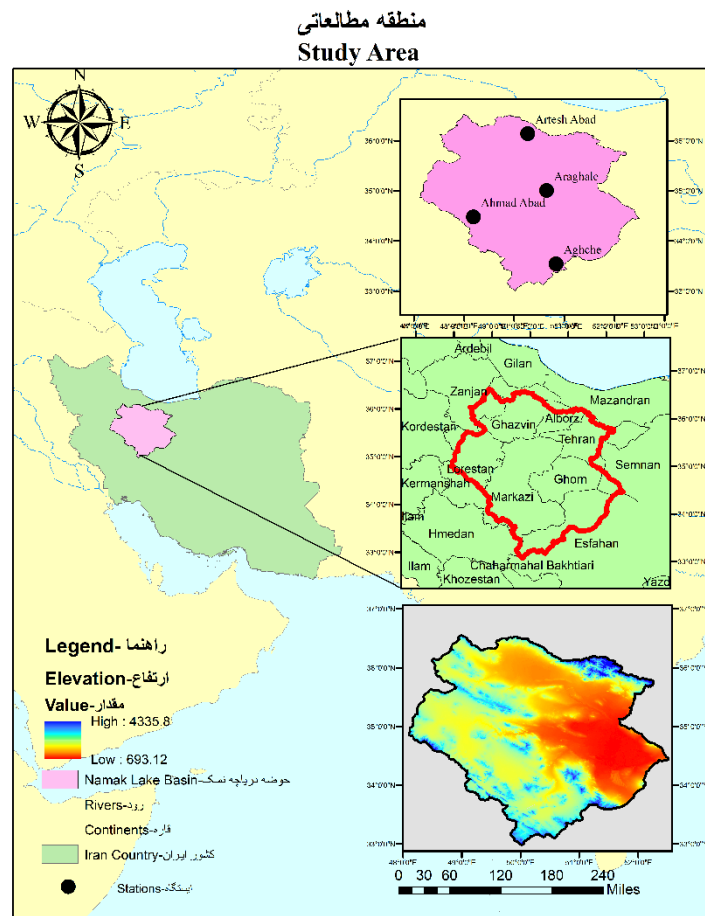
در ادامه این دسته، پژوهش‌هایی قرار می‌گیرند که بر تخمین و پیش‌بینی ارتفاع بارش در مناطق فاقد داده تمرکز دارند. در این راستا، حدادیان و یزدی (Haddadi & Yazdi, 2025) با استفاده از XGBoost^۱ نشان می‌دهند که این الگوریتم یادگیری ماشین برای مناطق فاقد اندازه‌گیری میدانی گزینه‌ای مناسب است و می‌تواند در کاربردهای مرتبط با مدیریت منابع آب به کار رود. بنابراین، این پژوهش را باید ذیل پیش‌بینی بارش در مناطق کم‌داده با رویکرد XGBoost طبقه‌بندی کرد. در نهایت، برخی مطالعات از مدل‌های آماری ریزمقیاس‌نمایی^۲ برای برآورد اثرات تغییر اقلیم بر بارش استفاده کرده‌اند. برای مثال، در مطالعه‌ای که از مدل SDSM^۳ در دو اقلیم خشک و فراخشک استفاده شد، نشان داده شد که؛ عملکرد مدل در منطقه خشک بهتر از فراخشک است و تا سال ۲۰۹۹، بارش سالانه در کرمان اندکی کاهش و در بم افزایش می‌یابد (Rezaei et al., 2014). این نتایج نشان می‌دهد که؛ مدل‌های آماری ریزمقیاس‌نما می‌توانند در برآورد اثرات تغییر اقلیم بر بارش محلی نقش مؤثری داشته باشند.

در کنار این‌ها، پژوهش‌هایی نیز رویکردهای ترکیب مدل‌ها^۴ یا مدل‌های هیبریدی را توسعه داده‌اند. یلدیز و همکاران (۲۰۲۳) با توسعه روش EK-stars^۵ مبتنی بر یادگیری گروهی، افزایش دقت پیش‌بینی بارش را گزارش کردند (Tuysuzoglu et al., 2023).

در جمع‌بندی پیشینه می‌توان بیان کرد که؛ اگرچه مطالعات متعددی از روش‌های یادگیری ماشین، یادگیری عمیق و مدل‌های ترکیبی برای پیش‌بینی بارش استفاده کرده‌اند، اما همچنان چند شکاف مهم وجود دارد: نخست آن که بررسی هم‌زمان عملکرد داده‌های CHIRPS در کنار طیف گسترده‌ای از مدل‌های یادگیری ماشین کمتر انجام شده و اغلب پژوهش‌ها بر داده‌های زمینی یا تعداد محدودی مدل تمرکز داشته‌اند؛ دوم آن که مقایسه نظام‌مند مدل‌های منفرد و ترکیبی در مقیاس ماهانه، به‌ویژه در مناطق خشک با رفتار بارشی نامنظم، کمتر مورد توجه قرار گرفته است؛ و سوم آن که اغلب مطالعات از وزندهی ساده برای ترکیب مدل‌ها استفاده کرده‌اند و رویکردهای بهینه‌سازی وزن‌ها، مانند روش BruteForce برای افزایش دقت پیش‌بینی، کمتر به کار رفته است. بنابراین مطالعه حاضر با ارزیابی جامع داده‌های CHIRPS، مقایسه طیفی از مدل‌های یادگیری ماشین و توسعه یک مدل ترکیبی با وزن‌های بهینه‌شده، به‌طور مستقیم این شکاف‌ها را هدف قرار می‌دهد.

- 1- Extreme Gradient Boosting
- 2- Statistical Downscaling Models
- 3- Statistical DownScaling Model
- 4- ensemble
- 5- Ensemble K-Star

6- Interquartile Range – IQR



شکل ۱- حوضه دریاچه نمک

Figure 1– Salt Lake Basin

به منظور آموزش و ارزیابی مدل‌ها، داده‌ها به دو بخش تقسیم گردید:

داده‌های آموزش: ژانویه ۲۰۰۰ تا دسامبر ۲۰۲۱

داده‌های آزمون: ژانویه ۲۰۲۲ تا دسامبر ۲۰۲۴

ارزیابی دقت داده‌های بارش CHIRPS

به منظور ارزیابی دقت داده‌های بارش CHIRPS، از داده‌های مشاهده‌ای چهار ایستگاه زمینی با پراکندگی مکانی مناسب در حوزه مورد مطالعه استفاده شد. ایستگاه‌های آغچه در استان اصفهان، ارتش‌آباد در استان قزوین، و ایستگاه‌های آراقلعه و احمدآباد در استان مرکزی طی دوره آماری ۱۰ ساله (۲۰۰۷ تا ۲۰۱۷) مورد بررسی قرار گرفتند. مشخصات مکانی ایستگاه‌ها در **جدول ۱** ارائه شده است. در این مطالعه، ابتدا داده‌های روزانه بار CHIRPS با داده‌های مشاهده‌ای هر ایستگاه به صورت جداگانه مقایسه شدند. برای این منظور، مقادیر بارش ماهواره‌ای متناظر با مختصات جغرافیایی هر ایستگاه استخراج و سپس شاخص‌های آماری ارزیابی گردید.

با توجه به یکنواخت بودن مقیاس داده‌ها در کل دوره زمانی نیازی به نرمال‌سازی داده‌ها وجود نداشت. در نتیجه داده‌ها بدون تغییر مقیاس وارد مرحله آموزش و آزمون مدل‌ها شدند. در این پژوهش، داده‌های بارش ماهانه به‌عنوان متغیر ورودی مدل‌ها مورد استفاده قرار گرفت. بدین منظور، میانگین بارش ماهانه استخراج‌شده از پایگاه داده CHIRPS (بر حسب mm) برای منطقه مورد مطالعه به‌عنوان یکی از ورودی‌های اصلی مدل‌ها در نظر گرفته شد. همچنین به‌منظور در نظر گرفتن تغییرات فصلی بارش در طول سال، شماره ماه نیز به‌عنوان یک متغیر ورودی به مدل‌ها اضافه شد. بنابراین ورودی مدل‌ها شامل دو متغیر اصلی یعنی میانگین بارش ماهانه و شماره ماه مربوطه بود. در این پژوهش از داده‌های بارش با زمان تأخیر^۱ استفاده نشده است و تمرکز اصلی بر بررسی رابطه بین مقادیر بارش ماهانه و تغییرات زمانی، آن در طول سال بوده است.

1- Lag Time

جدول ۱- مختصات ایستگاه‌های بررسی شده

Table 1- Coordinates of the investigated stations

نام ایستگاه Station name	استان Province	ارتفاع Elevation (m)	طول جغرافیایی Longitude (DMS)	عرض جغرافیایی Latitude (DMS)
آغچه (Aghcheh)	اصفهان (Esfahan)	2500	50-23-25	33-04-02
ارتش آباد (Artesh Abad)	قزوین (Ghazvin)	1750	49-25-47	35-39-31
آراقله (Ara Qaleh)	مرکزی (Markazi)	1641	49-43-19	35-06-05
احمدآباد (Ahmad Abad)	مرکزی (Markazi)	104	50-26-18	34-59-42

مدل‌های درختی و قاعده‌محور (M5P و M5Rules، Decision Tree، Random Forest، Bagging، Additive و Random SubSpace، Regression)، انجام شد. دلیل آزمون این طیف متنوع از مدل‌ها، بررسی توانایی آن‌ها در بازنمایی رفتار غیرخطی و نوسانی بارش ماهانه و همچنین ارزیابی این فرضیه بود؛ که مدل‌های ترکیبی و درختی به دلیل قابلیت مدیریت روابط پیچیده و کاهش واریانس، می‌توانند عملکرد بهتری نسبت به مدل‌های ساده‌تر ارائه دهند. این رویکرد امکان مقایسه جامع و انتخاب ساختار بهینه برای پیش‌بینی بارش در حوضه مورد مطالعه را فراهم می‌سازد. با استفاده از ابزار Grid Search در نرم‌افزار WEKA پارامترهای هر مدل متناسب با سری زمانی به صورت بهینه انتخاب شدند. در ادامه جدول ۲ توضیحات مدل‌ها را نشان می‌دهد.

ترکیب مدل‌ها

به منظور ارتقاء دقت و پایداری، پس از اجرای انفرادی مدل‌ها، ابتدا ترکیب‌های زوجی بر اساس میانگین‌گیری ساده از خروجی دو مدل ایجاد شدند. این ترکیب بدون اعمال وزن بوده و صرفاً به عنوان راهی برای کاهش خطا و افزایش اعتبار پذیری پیش‌بینی‌ها به کار رفت. در مجموع، ۱۰۵ ترکیب مختلف حاصل شد که ارزیابی و انتخاب مدل‌های برتر، مبتنی بر نتایج داده‌های آزمون صورت پذیرفت. پس از میانگین‌گیری ساده، از روش Brute force (جستجوی فراگیر برای وزن‌دهی بهینه) استفاده شده است. در این روش از صفر تا یک، به هر مدل با گام ۰/۰۱ وزن داده می‌شود؛ به طوری که مجموع وزن مدل‌ها، یک می‌شود و تمام حالت‌هایی که می‌شود به یک مدل وزن داده، امتحان می‌شود. در این روش تمام ترکیب‌های ممکن از وزن‌ها بررسی می‌شود تا بهترین حالت پیدا شود.

همچنین به منظور نمایش پراکنش مکانی ایستگاه‌ها در محدوده حوضه آبریز دریاچه نمک، موقعیت آن‌ها در نقشه منطقه مورد مطالعه نیز در شکل ۱ نمایش داده شده است. مدل‌های مورد استفاده در این مطالعه، ابتدا ۱۶ مدل مختلف یادگیری ماشین برای پیش‌بینی سری زمانی بارش مورد بررسی قرار گرفت. به دلیل عملکرد ضعیف یکی از مدل‌ها (مدل Random Tree با ضریب همبستگی کمتر از ۰/۵ و مقدار RMSE بالا)، مدل مزبور حذف گردید و نهایتاً ۱۵ مدل برتر انتخاب و ارزیابی شدند. مدل‌های مورد استفاده در این پژوهش طیف متنوعی از رویکردهای یادگیری ماشین شامل مدل‌های خطی، مبتنی بر فاصله، مبتنی بر هسته، درختی و مدل‌های ترکیبی را در بر می‌گیرند. هر یک از این الگوریتم‌ها دارای نقاط قوت و محدودیت‌های خاص خود هستند. به طور کلی، مدل‌های خطی مانند LR^1 دارای سادگی و تفسیرپذیری بالا هستند اما در بازنمایی روابط غیرخطی محدودیت دارند. مدل‌های مبتنی بر هسته مانند SVR و GP^2 توانایی مدل‌سازی روابط پیچیده را دارند ولی از نظر محاسباتی پرهزینه‌تر هستند. الگوریتم‌های درختی و مبتنی بر قواعد مانند M5Rules و M5P تعادل مناسبی میان دقت و تفسیرپذیری ایجاد می‌کنند. در نهایت، روش‌های ترکیبی مانند Bagging و جنگل تصادفی^۳ با کاهش واریانس و خطای مدل‌های پایه، پایداری پیش‌بینی را افزایش می‌دهند. انتخاب این مجموعه متنوع از مدل‌ها با هدف مقایسه جامع عملکرد و بررسی امکان بهبود دقت از طریق ترکیب آن‌ها صورت گرفته است. انتخاب این مدل‌ها با هدف پوشش رویکردهای مختلف مدل‌سازی شامل مدل‌های خطی (Linear Regression)، مبتنی بر فاصله (IBk)، مبتنی بر هسته (SMOreg و Gaussian Processes)، شبکه‌های عصبی (MLP و RBF)،

- 1- Linear Regression
- 2- Gaussian Processes
- 3- Random Forest

جدول ۲- توضیحات مدل‌های یادگیری ماشین مورد استفاده در پژوهش

Table 2- Description of the Machine Learning Models used in the study

مدل Model	معایب Disadvantages	مزایا Advantages	مرجع Reference
Gaussian Processes	هزینه محاسباتی بالا، حساس به انتخاب کرنل High computational cost; sensitive to kernel selection	توانایی مدل‌سازی روابط غیرخطی پیچیده؛ ارائه عدم قطعیت پیش‌بینی؛ عملکرد قوی در داده‌های کوچک Ability to model complex nonlinear relationships; provides prediction uncertainty; strong performance with small datasets	(Schulz et al., 2018)
Random Forest	کاهش تفسیرپذیری؛ نیاز به تنظیم تعداد درخت‌ها و پارامترهای دیگر Reduced interpretability; requires tuning of tree numbers and parameters	مقاوم به نویز؛ کاهش بیش‌برازش نسبت به درخت منفرد؛ عملکرد قوی در داده‌های غیرخطی Robust to noise; reduces overfitting compared to a single tree; strong nonlinear performance	(Biau & Scornet 2016)
MLP Regressor	نیاز به تنظیم زیاد پارامترها؛ حساس به داده کم Requires extensive hyperparameter tuning; sensitive to small datasets	قدرت بالا در مدل‌سازی روابط غیرخطی پیچیده؛ مناسب سری‌های زمانی Strong capability for modeling complex nonlinear patterns; suitable for time series	(Goodfello Bengio & Courville, 2016)
Linear Regression	ناتوان در مدل‌سازی روابط غیرخطی Cannot model nonlinear relationships	ساده، سریع، قابل تفسیر؛ مبنای مقایسه Simple, fast, interpretable; baseline method	(James et al., 2021)
SMOreg	انتخاب پارامتر کرنل پیچیده Complex kernel parameter selection	عملکرد پایدار در داده‌های نویزی؛ کنترل بیش‌برازش Stable performance on noisy data; effective control of overfitting	(Cervantes et al., 2020)
IBk	حساس به مقیاس داده؛ عملکرد ضعیف در داده پرتعداد Sensitive to data scaling; weak performance on large datasets	ساده؛ بدون فرض توزیع؛ مناسب داده‌های محلی Simple; no distribution assumptions; suitable for local patterns	(Taunk et al., 2019)
LWL	وزن High computational cost during prediction; sensitive to weighting parameters	مدل‌سازی رفتار موضعی؛ انعطاف‌پذیر Models local behavior; highly flexible	(Abinaya, 2020)
Additive Regression	خطر بیش‌برازش در تکرار زیاد Risk of overfitting with too many iterations	کاهش بایاس؛ بهبود مدل پایه Reduces bias; improves the base learner	(Chen & Guestrin, 2016)
Bagging	کاهش تفسیرپذیری Reduced interpretability	کاهش واریانس؛ پایداری بالا Reduces variance; high stability	(Sagi & Rokach, 2018)
Random Committee	وابسته به کیفیت مدل پایه Dependent on quality of base learner	بهبود پایداری مدل پایه؛ ساده در اجرا Improves stability of the base model; easy to implement	(Dong et al., 2020)
Random SubSpace	احتمال حذف ویژگی‌های مهم Reduces correlation among models; suitable for high-dimensional data	کاهش همبستگی بین مدل‌ها؛ مناسب داده‌های با ویژگی زیاد Possible removal of important features	(Kuncheva, 2014)
Decision Table	عملکرد ضعیف در روابط پیچیده Poor performance on complex relationships	تفسیرپذیری بالا؛ سادگی High interpretability; simple	(Fürnkranz et al., 2012)
M5Rules	حساس به داده‌های نویزی Sensitive to noisy data	ترکیب درخت تصمیم و رگرسیون خطی؛ تفسیرپذیر و دقیق Combines decision trees with linear regression; interpretable and accurate	(Holmes et al., 1999)
M5P	امکان بیش‌برازش در داده کم Possible overfitting with small datasets	مناسب روابط غیرخطی؛ ساختار قابل تفسیر Suitable for nonlinear relationships; interpretable	(Wang & Witten, 1997)
RBF Regressor	حساس به انتخاب مراکز و پارامتر Sensitive to selection of centers and parameters	همگرایی سریع‌تر از MLP؛ مناسب روابط موضعی Faster convergence than MLP; suitable for local relationships	(Wu et al., 2012)

معیارهای ارزیابی عملکرد مدل‌ها

به‌منظور سنجش و مقایسه دقیق عملکرد مدل‌های منفرد و ترکیبی، از شاخص‌های آماری و معیارهای خطای زیر استفاده شد:

ضریب همبستگی (Correlation): این شاخص میزان همبستگی و ارتباط خطی بین مقادیر مشاهده‌شده و مقادیر پیش‌بینی‌شده توسط مدل را نشان می‌دهد و مقدار آن بین -۱ تا +۱ تغییر می‌کند. هرچه مقدار این شاخص به ۱ نزدیک‌تر باشد، نشان‌دهنده تطابق بهتر بین داده‌های مشاهده‌ای و مقادیر برآورد شده توسط مدل است. معمولاً از توان دوم این ضریب استفاده می‌شود.

$$R^2 = \frac{[\sum_{i=1}^n (O_i - \bar{O}) \cdot (S_i - \bar{S})]^2}{\sum_{i=1}^n (O_i - \bar{O})^2 \cdot \sum_{i=1}^n (S_i - \bar{S})^2} \quad (۱)$$

ریشه میانگین مربعات خطا (RMSE^۱): این شاخص یکی از متداول‌ترین معیارها برای ارزیابی دقت مدل‌ها است و میزان انحراف مقادیر پیش‌بینی‌شده از مقادیر مشاهده‌شده را نشان می‌دهد. مقدار کمتر RMSE نشان‌دهنده عملکرد بهتر مدل است.

این شاخص در بازه ۰ تا +۱ است و هرچه مقدار این شاخص به ۰ نزدیک‌تر باشد نشان‌دهنده دقت بالاتر است.

$$RMSE = \sqrt{\frac{\sum_{i=1}^n (O_i - S_i)^2}{n}} \quad (۲)$$

میانگین قدر مطلق خطا (MAE^۲): این شاخص میانگین قدر مطلق اختلاف بین مقادیر مشاهده‌شده و پیش‌بینی‌شده را نشان می‌دهد و نسبت به خطاهای بزرگ حساسیت کمتری نسبت به RMSE دارد. مقدار کمتر MAE بیانگر دقت بیشتر مدل است. این شاخص در بازه ۰ تا +۱ است و هرچه مقدار این شاخص به ۰ نزدیک‌تر باشد نشان‌دهنده دقت بالاتر است.

$$MAE = \frac{\sum_{i=1}^n (O_i - S_i)}{n} \quad (۳)$$

میانگین توان دو خطا (MSE^۳): این شاخص میانگین مربع اختلاف بین مقادیر مشاهده‌شده و پیش‌بینی‌شده است و به‌دلیل مجذور شدن خطاها، نسبت به خطاهای بزرگ حساسیت بیشتری دارد. این شاخص در بازه ۰ تا +۱ است و هرچه مقدار این شاخص به ۰ نزدیک‌تر باشد نشان‌دهنده دقت بالاتر است.

$$MSE = \frac{\sum_{i=1}^n (O_i - S_i)^2}{n} \quad (۴)$$

بایاس (Bias): شاخص بایاس میزان تمایل مدل به بیش‌برآوردی یا کم‌برآوردی مقادیر را نشان می‌دهد. مقدار مثبت بیانگر بیش‌برآوردی و مقدار منفی نشان‌دهنده کم‌برآوردی مدل است و

هرچه به ۰ نزدیک‌تر باشد نزدیک‌تر بودن به مقادیر واقعی را نشان می‌دهد.

$$SE = \bar{O} - \bar{S} \quad (۵)$$

شاخص نش-ساتکلیف (NSE^۴): این شاخص یکی از معیارهای پرکاربرد در ارزیابی مدل‌های هیدرولوژیکی است و میزان کارایی مدل در بازتولید داده‌های مشاهده‌ای را نشان می‌دهد. مقدار این شاخص بین منفی بی‌نهایت تا ۱ تغییر می‌کند و هرچه مقدار آن به ۱ نزدیک‌تر باشد، نشان‌دهنده عملکرد بهتر مدل است.

$$NSE = \frac{\sum_{i=1}^n (S_i - O_i)^2}{\sum_{i=1}^n (O_i - \bar{O})^2} \quad (۶)$$

O_i مقدار مشاهده شده در گام زمانی i

$\bar{O}$ میانگین مقدار مشاهداتی

$\bar{S}$ میانگین مقادیر پیش‌بینی شده

S_i مقدار پیش‌بینی شدت در گام زمانی i

n تعداد نمونه یا تعداد داده

کلیه این معیارها برای تمام مدل‌های منفرد و ترکیبی محاسبه شد تا بهترین ساختار مدل‌سازی بارش ماهانه در محدوده مطالعه انتخاب گردد.

نتایج و بحث

نتایج ارزیابی دقت ماهواره

پس از انتخاب ایستگاه‌های باران‌سنجی در محدوده حوضه آبریز دریاچه نمک ایستگاه‌های آغچه در استان اصفهان، ارتش‌آباد در استان قزوین، و ایستگاه‌های آراقلعه و احمدآباد در استان مرکزی طی دوره آماری ۱۰ ساله (۲۰۰۷ تا ۲۰۱۷) مورد بررسی قرار گرفتند و داده‌های بارش ماهانه این ایستگاه‌ها گردآوری شد و برای ارزیابی دقت داده‌های ماهواره‌ای CHIRPS مورد استفاده قرار گرفت. مقایسه بین داده‌های زمینی و داده‌های متناظر CHIRPS در مقیاس ماهانه انجام شد که نتایج آن در **جدول ۳** ارائه شده است. نتایج ارائه‌شده در **جدول ۳** بیانگر میانگین شاخص‌های آماری محاسبه‌شده برای کل ایستگاه‌ها است. به‌عبارت دیگر، پس از محاسبه شاخص‌های ارزیابی برای هر ایستگاه به‌صورت مجزا، مقادیر نهایی **جدول ۳** از طریق میانگین‌گیری از نتایج چهار ایستگاه به‌دست آمده‌اند تا عملکرد کلی داده‌های CHIRPS در سطح حوضه مورد مطالعه ارزیابی شود.

1- RMSE: Root Mean Square Error

2- MAE: Mean Absolute Error

3- MSE: Mean Square Error

4- NSE: Nash-Sutcliffe Model Efficiency

جدول ۳- نتایج ارزیابی دقت داده‌های CHIRPS در منطقه مورد مطالعه

Table 3- Accuracy assessment results of CHIRPS data in the study area

منبع داده Data Source	نش-سایتیف NSE (%)	میانگین قدر مطلق خطا MAE (%)	میانگین توان دو خطا MSE (%)	بایاس BAIAS (mm)	همبستگی Correlation (-)
CHIRPS	70.00	31.36	19.42	-00.09	0.71

پایین بودن قدرمطلق مقدار بایاس به معنای دقت بالاتر مدل در برآورد میانگین واقعی مشاهدات است. بنا بر نتایج، مدل Gaussian Processes کمترین مقدار بایاس را دارد (تقریباً ۰/۰۴) که بیانگر نزدیکی بیشتر نتایج پیش‌بینی‌شده توسط این مدل به میانگین مقادیر واقعی است. در مقابل، مدل Random SubSpace بیشترین مقدار بایاس را نشان می‌دهد (۰/۱۹) و بیانگر انحراف سیستماتیک بیشتر در پیش‌بینی‌های این مدل می‌باشد. بایاس مثبت به معنای بیش‌برآورد و بایاس منفی به معنای کم‌برآورد است. در شکل ۲ (دیاگرام تیلور) برای مقایسه واضح‌تر مدل‌ها آورده شده است.

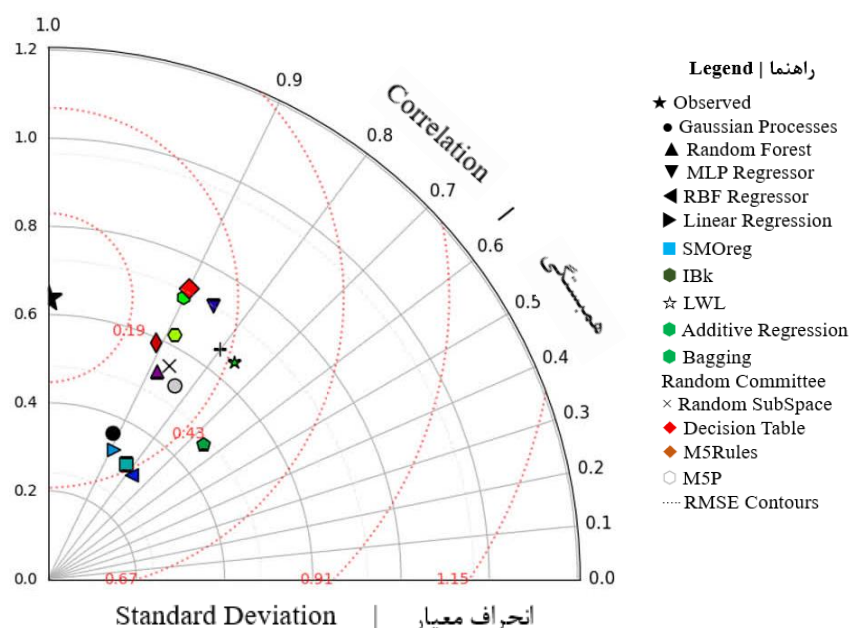
نمودار تیلور به‌منظور ارزیابی همزمان عملکرد مدل‌ها از نظر انحراف معیار، ضریب همبستگی و خطای میانگین مربعات خطا نسبت به داده‌های مشاهداتی به‌کار گرفته شد. در این نمودار، نقطه مرجع (Observed) نشان‌دهنده داده‌های مشاهداتی است که دارای ضریب همبستگی برابر با یک و مقدار میانگین مربعات خطا برابر با صفر می‌باشد. فاصله هر مدل از این نقطه، بیانگر میزان خطای آن نسبت به مشاهدات است. محور افقی در نمودار تیلور بیانگر انحراف معیار است که نشان‌دهنده دامنه تغییرات متغیر مورد بررسی می‌باشد. این پارامتر یک ویژگی آماری از سری زمانی بوده و نباید با شاخص‌های خطای مدل نظیر میانگین مربعات خطا اشتباه گرفته شود. ضریب همبستگی در نمودار تیلور به‌صورت زاویه‌ای نمایش داده می‌شود، به‌طوری که زاویه بین هر نقطه مدل و محور افقی متناظر با مقدار ضریب همبستگی آن مدل با داده‌های مشاهداتی است. مقادیر بزرگ‌تر ضریب همبستگی نشان‌دهنده تطابق بهتر الگوی تغییرات مدل با مشاهدات می‌باشند. در این نمودار، داده‌های مشاهداتی برابر با سری زمانی بارش ماهانه استخراج‌شده از پایگاه داده CHIRPS بوده است که به‌عنوان مرجع ارزیابی عملکرد مدل‌ها در نظر گرفته شد. لازم به ذکر است که این ارزیابی برای دوره آزمون مطالعه، شامل دو سال پایانی داده‌ها، انجام شده است. با توجه به نمودار، مدل M5Rules نزدیک‌ترین موقعیت را به نقطه‌ی مرجع دارد، که نشان‌دهنده دقت بالا، همبستگی قوی و انطباق مناسب با پراکندگی داده‌های واقعی است.

براساس این نتایج، مقدار R^2 برابر ۰/۷۱ نشان‌دهنده توان نسبتاً مناسب CHIRPS در توضیح تغییرات زمانی بارش در ایستگاه‌ها است. مقدار NSE حدود ۷۰ درصد نیز عملکرد قابل قبول این مجموعه داده در بازتولید الگوی بارش ماهانه را تأیید می‌کند. همچنین مقادیر MAE و MSE به‌ترتیب برابر با ۳۱/۳۶ و ۱۹/۴۲ درصد نشان می‌دهند که اگرچه همخوانی کلی مناسب است، اما بخشی از خطا باقی می‌ماند. مقدار BIAS منفی (۰/۰۹-) نیز حاکی از آن است که CHIRPS در مجموع تمایل اندکی به کم‌برآوردی بارش نسبت به داده‌های ایستگاهی دارد. با این حال، با توجه به گستردگی منطقه، پراکندگی ایستگاه‌ها و شرایط خشک و نیمه‌خشک حوضه، دقت به‌دست‌آمده قابل قبول بوده و استفاده از داده‌های CHIRPS برای تحلیل‌های بعدی در حوضه آبریز توجیه‌پذیر است.

عملکرد مدل‌های منفرد

در این بخش، عملکرد ۱۵ مدل یادگیری ماشین به‌صورت جداگانه برای پیش‌بینی بارش ماهانه در منطقه مورد مطالعه (حوضه آبریز دریاچه نمک) بررسی شده است. ارزیابی مدل‌ها بر اساس داده‌های آزمون (ژانویه ۲۰۲۲ تا دسامبر ۲۰۲۴) انجام شده و از شاخص‌هایی همچون ضریب همبستگی، ریشه میانگین مربع خطا (RMSE)، میانگین قدرمطلق خطا (MAE) استفاده گردیده است. در ادامه در جدول ۴ نتایج پیش‌بینی مدل‌ها آورده شده است.

نتایج جدول ۴ نشان می‌دهد که؛ عملکرد مدل‌ها نه تنها به میزان خطای کلی (MAE و MSE) بلکه به ترکیب بایاس سیستماتیک و پراکندگی خطاها وابسته است. مدل‌هایی مانند M5Rules و Additive Regression دارای بایاس کم، پراکندگی پایین و میانه قدرمطلق خطای کوچک هستند که نشان‌دهنده پیش‌بینی دقیق و پایدار است. در مقابل، مدل‌هایی مانند IBk با پراکندگی بالای خطا و بایاس قابل توجه، علی‌رغم میانه خطای نسبتاً پایین، برخی پیش‌بینی‌های پرت داشته و MSE بالاتری ثبت کرده‌اند. میانه قدرمطلق خطاها شاخص ناپارامتری خطا است که حساسیت کمتری به نقاط پرت دارد و نشان می‌دهد که خطای غالب مدل‌ها در بیشتر نمونه‌ها چقدر است. بنابراین این مجموع شاخص‌ها تصویر کامل‌تر و دقیق‌تری از عملکرد مدل‌ها ارائه می‌دهد.



شکل ۲- نمودار تیلور عملکرد مدل‌ها در پیش‌بینی بارش ماهانه طی دوره آزمون

Figure 2- Taylor diagram of the performance of the models monthly precipitation prediction during the test period

می‌شود، عملکرد قابل قبولی دارند، اما در مواجهه با ماه‌های با رفتار غیرخطی یا پراکندگی شدید داده‌ها عملکرد ضعیف‌تری نشان می‌دهند. در مجموع، بهتر عمل کردن برخی مدل‌ها نتیجه ترکیبی از توانایی آن‌ها در مدیریت غیرخطی بودن داده‌ها، میزان حساسیت به نقاط پرت، قابلیت تعمیم‌پذیری و نحوه یادگیری الگوهای زمانی است. این تحلیل کمک می‌کند تا تفاوت مدل‌ها از دیدگاه علمی و بر اساس ماهیت آماری داده‌های بارش توضیح داده شود.

ترکیب مدل‌ها با وزن‌دهی ساده (میانگین ساده)

در مدل‌های ترکیبی، هر الگوریتم می‌تواند بخشی از فرآیند یادگیری و پیش‌بینی را بر اساس ویژگی‌های خاص خود انجام دهد. به عبارت دیگر، الگوریتم‌های مختلف دارای توانایی‌های متفاوتی در استخراج الگوهای داده هستند؛ به‌طوری‌که برخی مدل‌ها در شناسایی روابط غیرخطی پیچیده عملکرد بهتری دارند، در حالی که برخی دیگر در تعمیم‌پذیری و کاهش خطای پیش‌بینی مؤثرتر عمل می‌کنند. ترکیب این مدل‌ها می‌تواند با بهره‌گیری همزمان از نقاط قوت آن‌ها موجب بهبود عملکرد نهایی شود. در این پژوهش ترکیب‌های مختلفی از مدل‌ها مورد بررسی قرار گرفت تا مشخص شود کدام ترکیب قادر است خطای پیش‌بینی را کاهش داده و دقت مدل را افزایش دهد. با توجه به تعداد مدل‌های مورد استفاده، در مجموع ۱۰۵ حالت مختلف از ترکیب مدل‌ها ایجاد شد که بررسی همه آن‌ها در متن مقاله

پس از آن، مدل‌های Additive Regression و Random Forest نیز عملکرد قابل قبولی داشته‌اند. در مقابل، مدل‌هایی مانند IBk و RBF Regressor در فاصله دورتری از نقطه مرجع قرار دارند که نشان‌دهنده خطای بیشتر و همبستگی ضعیف‌تر آن‌ها است. نتایج حاصل از ارزیابی مدل‌ها نشان می‌دهد که تفاوت در عملکرد آن‌ها ناشی از ویژگی‌های الگوریتمی و نحوه یادگیری الگوهای موجود در داده‌های بارش است. مدل‌هایی که ساختار مبتنی بر درخت تصمیم دارند (نظیر جنگل تصادفی) معمولاً توانایی بالایی در شناسایی روابط غیرخطی و تعامل میان متغیرها دارند و در داده‌هایی با نوسان‌های شدید و رفتارهای غیرخطی، عملکرد پایدارتر و مقاوم‌تری نشان می‌دهند. این ویژگی سبب شده تا این مدل‌ها در بسیاری از ماه‌ها خطای کمتری نسبت به سایر روش‌ها داشته باشند. از سوی دیگر، مدل‌های شبکه عصبی مانند MLP قادرند الگوهای پیچیده‌تری را در داده‌ها بیاموزند، اما عملکرد آن‌ها به مقدار داده، ساختار شبکه و تنظیم پارامترها بسیار حساس است. در داده‌های بارش ماهانه که از نظر حجم محدود بوده و رفتارهای ناگهانی و غیرایستا دارند، حساسیت این مدل‌ها می‌تواند منجر به نوسان در دقت پیش‌بینی شود. با این حال، در ماه‌هایی که الگوی بارش دارای ساختار قابل یادگیری است، مدل MLP نتایج بسیار خوبی ارائه می‌دهد. مدل‌های خطی یا ساده‌تر در شرایطی که رابطه میان ورودی و خروجی کمتر پیچیده است یا تغییرات نسبتاً منظم‌تری مشاهده

نهایت ۱۲ مدل ترکیبی برتر از میان ۱۰۵ حالت که کارایی بالایی داشته‌اند آورده شده است؛ که در **جدول ۵** نمایش داده شده است.

جدول ۵ نشان می‌دهد که؛ مقایسه عملکرد مدل‌های منفرد و بهترین ترکیب هر مدل را از نظر شاخص نش-ساتکلیف (NSE) نشان می‌دهد. همان‌طور که مشخص است؛ در تمامی موارد، مقدار NSE برای ترکیب مدل‌ها بالاتر از حالت منفرد است. این موضوع بیانگر آن است که؛ ترکیب ساده دو مدل می‌تواند به‌طور معناداری دقت پیش‌بینی را افزایش دهد. بیشترین بهبود مربوط به مدل‌هایی است؛ که در حالت منفرد عملکرد ضعیف‌تری داشته‌اند، از جمله IBk که به تنهایی NSE ۳۴/۳ را دارد و پس از ترکیب شدن با مدل Decision Table عملکرد آنها به ۷۴/۴ می‌رسد و RBF Regressor که به تنهایی NSE ۴۷/۷ را داشته‌اند با ترکیب شدن با مدل Decision Table، NSE به ۸۲/۴ دست می‌یابند در همچنین، مدل‌هایی در حالت منفرد نیز عملکرد خوبی داشتند، در ترکیب نیز دقت بالاتری حفظ کرده‌اند مانند Additive Regression با NSE ۷۵/۲ پس از ترکیب شدن با مدل M5Rules به عملکرد ۸۷/۰ می‌رسد (M5Rules نیز به تنهایی عملکرد ۸۰/۰ را داشته است). این یافته‌ها اثبات می‌کند که رویکرد ترکیب دو به دو، نه تنها ضعف برخی مدل‌ها را جبران می‌کند، بلکه می‌تواند پایداری دقت پیش‌بینی را نیز افزایش دهد.

MAE برای مدل‌های ترکیبی نشان می‌دهد؛ که ترکیب Additive Regression - M5Rules با مقدار MAE برابر با ۱۵/۵، کمترین میزان خطای مطلق را دارد. این مسئله بیانگر آن است که این مدل، در پیش‌بینی بارش‌ها، به‌طور میانگین کمترین انحراف از مقادیر واقعی را داشته است. مدل‌های Decision Table - M5Rules و Additive Regression - M5P نیز با MAE های نزدیک به این مقدار، در رتبه‌های دوم و سوم قرار دارند

امکان‌پذیر نیست؛ از این‌رو تنها نتایج مهم‌ترین و کاراترین ترکیب‌ها ارائه شده است. لازم به ذکر است که؛ در روش‌های ترکیب مدل‌ها، انتخاب مدل‌ها صرفاً بر اساس بهترین عملکرد منفرد آنها انجام نمی‌شود. در بسیاری از موارد، مدل‌هایی که به‌صورت جداگانه دقت بالایی دارند ممکن است الگوی خطای مشابهی داشته باشند و در نتیجه ترکیب آن‌ها به بهبود قابل توجهی در نتایج منجر نشود. در مقابل، ترکیب مدل‌هایی که دارای الگوهای خطای متفاوت یا مکمل هستند می‌تواند منجر به کاهش خطای کلی شود. به بیان دیگر، در شرایطی که یک مدل در برخی زمان‌ها دچار بیش‌برآورد و مدل دیگر کم‌برآورد باشد، ترکیب خروجی آن‌ها می‌تواند این خطاها را تا حدی خنثی کرده و عملکرد بهتری نسبت به هر مدل به‌صورت منفرد ارائه دهد. بر همین اساس، در این پژوهش ترکیب‌های مختلفی از مدل‌ها مورد بررسی قرار گرفت تا ترکیبی که بیشترین توانایی در کاهش خطا و بهبود دقت پیش‌بینی دارد شناسایی شود.

در این مرحله، به‌منظور افزایش دقت پیش‌بینی، مدل‌های منتخب به‌صورت دو به دو با یکدیگر ترکیب شدند. در هر ترکیب، مقدار نهایی پیش‌بینی‌شده، میانگین حسابی خروجی دو مدل پایه بود. این رویکرد منجر به تولید ۱۰۵ حالت ترکیبی مختلف شد. میانگین حسابی یکی از ساده‌ترین و در عین حال پرکاربردترین روش‌ها در ترکیب مدل‌ها محسوب می‌شود و به‌دلیل توانایی آن در کاهش اثر خطاهای منفرد مدل‌ها، می‌تواند منجر به بهبود پایداری و دقت پیش‌بینی شود. در این روش، مقدار پیش‌بینی نهایی از طریق میانگین‌گیری از خروجی مدل‌های منتخب به‌دست می‌آید.

هدف از این بخش بررسی این موضوع بود که؛ آیا ترکیب مدل‌ها می‌تواند عملکرد بهتری نسبت به مدل‌های منفرد ارائه دهد یا خیر؟ برای این منظور، شاخص‌های ارزیابی ضریب همبستگی، MAE، MSE، NSE، BIAS برای تمام ترکیب‌ها محاسبه گردید. در

جدول ۵- معیارهای ارزیابی ۱۲ ترکیب برتر مدل‌ها

Table 5- Evaluation metrics of the top 12 model combinations

مدل MODEL	میانگین خطای مطلق MAE (%)	میانگین مربعات خطا MSE (%)	شاخص نش-ساتکلیف NSE (-)	همبستگی Correlation (-)	بایاس BAIAS (mm)
Gaussian Processes -Additive Regression	18.41	06.00	0.846	0.878	-0.119
Gaussian Processes -Decision Table	21.91	06.91	0.824	0.892	-0.213
Gaussian Processes -M5Rules	23.30	07.51	0.809	0.939	-0.135
Random Forest-Additive Regression	18.10	06.90	0.825	0.885	-0.187
Random Forest-M5Rules	21.61	07.71	0.803	0.944	-0.203
RBF Regressor-Decision Table	21.52	06.63	0.831	0.680	-0.049
Linear Regression -Additive Regression	22.18	07.59	0.808	0.852	-0.235
SMOreg -Decision Table	20.63	06.31	0.838	0.755	-0.039
LWL-M5Rules	19.76	07.81	0.801	0.685	-0.173
Additive Regression-M5Rules	15.55	05.08	0.873	0.832	-0.015
Additive Regression-M5P	0.18.02	06.20	0.841	0.768	-0.020
Decision Table-M5Rules	17.80	06.61	0.831	0.856	-0.272

این مدل‌ها توانسته‌اند با حفظ دقت بالا، پیش‌بینی‌هایی کم‌خطا در اغلب بازه‌های زمانی ارائه دهند. در مقابل، ترکیب Gaussian Processes - M5Rules با MAE برابر با ۰/۲۳۳، بیشترین خطای مطلق را داشته و عملکرد نسبتاً ضعیف‌تری در این شاخص نشان داده است. شاخص MSE که به خطاهای بزرگ‌تر حساس‌تر است، روندی مشابه با MAE را نشان می‌دهد. مدل Additive Regression - M5Rules با MSE برابر با ۰/۰۵ بهترین عملکرد را از نظر کمترین میزان خطای مربعی دارد. مدل‌های Additive Regression - M5P و Gaussian Processes - Additive Regression نیز با مقادیر نزدیک، عملکرد نسبتاً خوبی از خود نشان داده‌اند. در این شاخص نیز، مدل Gaussian Processes - M5Rules با MSE برابر با ۰/۰۷۵ در رتبه پایین‌تری قرار گرفته و بیانگر آن است که برخی از پیش‌بینی‌های این مدل دارای خطاهای بزرگ‌تری بوده‌اند. عملکرد ضعیف‌تر مدل Gaussian Processes - M5Rules در شاخص MSE به این دلیل است که این مدل در برخی نمونه‌ها دچار خطاهای نسبتاً بزرگ‌تری شده است. ترکیب ساختاری این مدل، که شامل قواعد M5Rules و فرآیندهای گاوسی است. در نتیجه، اگرچه عملکرد کلی مدل مناسب است، اما وجود چند پیش‌بینی با انحراف زیاد باعث افزایش MSE و قرار گرفتن آن در رتبه پایین‌تر شده است. نمودار ضریب همبستگی نشان می‌دهد که؛ مدل‌های ترکیبی مبتنی بر درختان تصمیم به‌ویژه ترکیب Random Forest - M5Rules بالاترین میزان تطابق زمانی را با داده‌های واقعی داشته‌اند. این مدل‌ها قادرند نوسانات غیرخطی و تغییرات ناگهانی بارش را بهتر دنبال کنند، زیرا ساختار درختی آن‌ها الگوهای شاخه‌ای و تغییرات رژیمی را به‌خوبی بازنمایی می‌کند. در مقابل، ترکیب‌هایی که شامل مدل‌های مبتنی بر توابع شعاعی یا همسایگی بوده‌اند (مانند RBF Regressor - Decision Table)، همبستگی پایین‌تری با سری زمانی مشاهده‌شده داشته‌اند. دلیل این موضوع، ویژگی ذاتی این مدل‌ها در هموارسازی خروجی است؛ این دسته مدل‌ها تمایل دارند تغییرات تند و نوسانات ریزمقیاس بارش را تعدیل کنند، بنابراین تطابق زمانی آن‌ها با رفتار واقعی بارش کاهش می‌یابد. مدل‌های درختی الگوهای غیرخطی و نقاط شکست را بهتر تشخیص می‌دهند، در حالی که مدل‌های مبتنی بر توابع پایه رفتار هموارتری ارائه می‌کنند و نوسانات واقعی را کمتر منعکس می‌سازند.

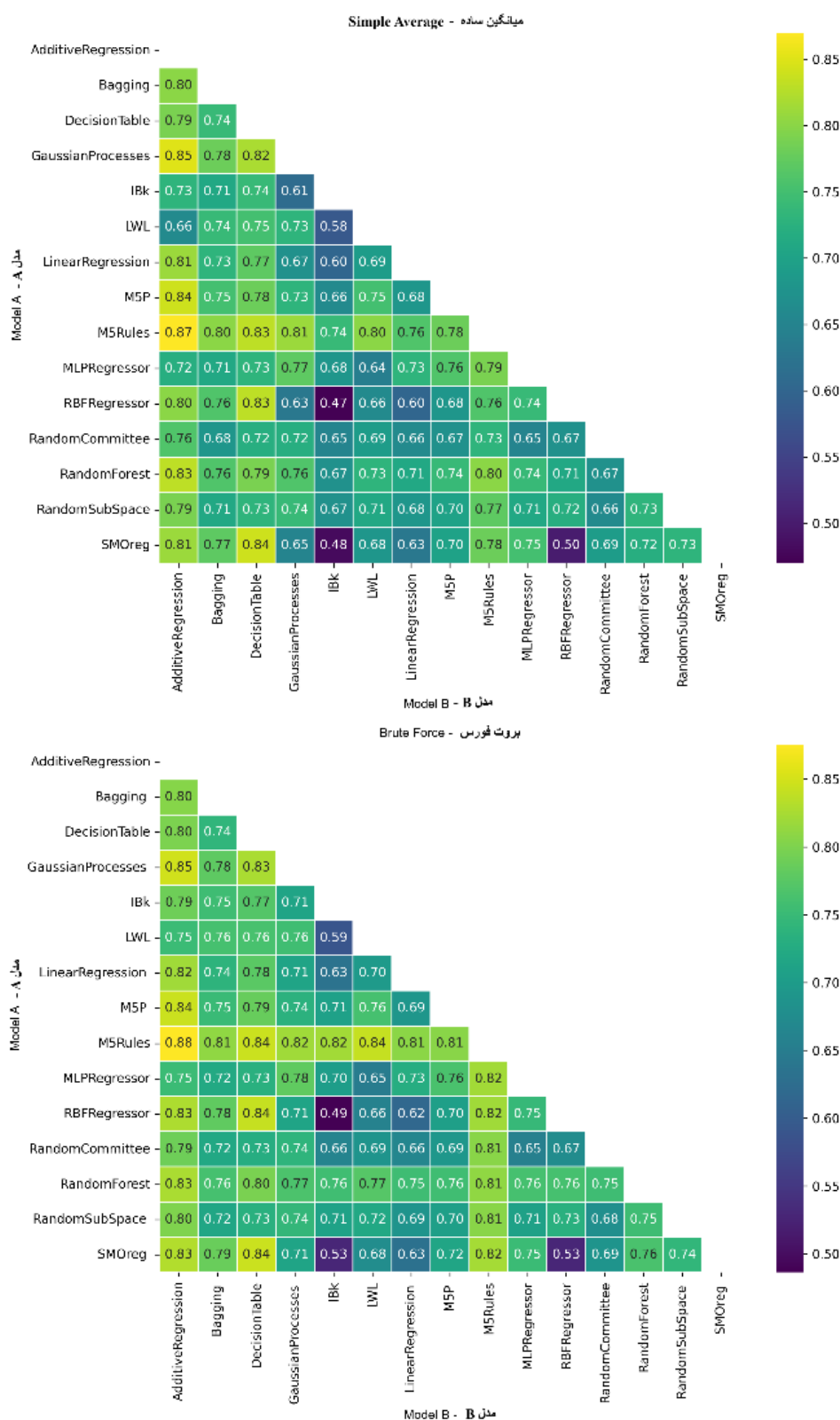
وزن‌دهی به روش BruteForce

در این روش از صفر تا یک به هر مدل با گام ۰/۰۱ وزن داده می‌شود؛ به‌طوری‌که مجموع وزن مدل‌ها، یک می‌شود. تمام حالت‌های تحلیل مقایسه‌ای ارائه شده در زیر، تفاوت‌های عملکردی بین دو رویکرد میانگین‌گیری ساده (Simple Average) و

BruteForce را بر اساس شاخص ضریب نش-ساتکلیف بررسی می‌کند. در این روش تمام ترکیب‌های ممکن از وزن‌ها بررسی می‌شود تا بهترین حالت پیدا شود. شکل ۳ نتایج وزن‌دهی ساده در مقابل وزن‌دهی به روش BruteForce را نشان می‌دهد. بررسی عملکردی در شکل ۳، تفاوت کارایی بین دو استراتژی ترکیب مدل‌ها، یعنی روش میانگین‌گیری ساده (Simple Average) و روش وزن‌دهی بهینه (BruteForce) را به خوبی آشکار می‌سازد. در نگاه کلی، مشاهده می‌شود که روش BruteForce با اختصاص وزن‌های بهینه به ترکیب‌های دوگانه مدل‌ها، توانسته است مقادیر شاخص NSE را در اکثر خانه‌های ماتریس بهبود ببخشد. با مقایسه دقیق‌تر دو نمودار، مشخص است که؛ در روش میانگین‌گیری ساده، رنگ‌های تیره‌تر (بنفش و سرمه‌ای) که نشان‌دهنده NSE پایین هستند، به‌ویژه در ترکیب‌هایی که شامل مدل‌های ضعیف‌تری همچون RBF Regressor و IBk می‌شوند، بیشتر دیده می‌شود. این موضوع نشان می‌دهد که در روش میانگین‌گیری ساده، ضعف‌های یک مدل به‌طور مستقیم بر خروجی نهایی ترکیب اثر منفی می‌گذارد. در مقابل، در نمودار مربوط به رویکرد BruteForce، بخش بزرگی از ماتریس به طیف رنگی زرد و سبز روشن متمایل شده است که نشان‌دهنده ارتقای سطح دقت به محدوده NSE بالای ۰/۸۰ است. به‌عنوان مثال، ترکیب مدل M5Rules با سایر الگوریتم‌ها در روش Brute Force به‌طور مداوم مقادیری NSE بالای ۰/۸۴ را ثبت کرده است، در حالی‌که در روش میانگین‌گیری ساده، این مقادیر نوسان بیشتری دارند. این برتری به این دلیل است که الگوریتم BruteForce با جستجوی سیستماتیک، وزن مدل‌های دارای خطای بالا را کاهش داده و بر وزن مدل‌های دقیق‌تر در ترکیب نهایی می‌افزاید. در مجموع، نتایج نشان می‌دهد که اگرچه میانگین‌گیری ساده روشی سریع و بدون پیچیدگی محاسباتی است، اما استفاده از روش‌های وزن‌دهی هوشمند مانند BruteForce می‌تواند به‌طور قابل توجهی خطاهای هم‌پوشان را خنثی کرده و عملکرد سیستم‌های ترکیبی را به سطحی بالاتر از تک‌تک مدل‌های پایه ارتقا دهد.

جمع‌بندی تحلیل نتایج

تحلیل جامع نتایج به‌دست‌آمده از مدل‌های منفرد و ترکیبی نشان می‌دهد که؛ دقت پیش‌بینی بارش ماهانه نه‌تنها به نوع الگوریتم مورد استفاده، بلکه به نحوه تعامل و هم‌پوشانی خطاهای مدل‌ها وابسته است. نتایج مدل‌های منفرد حاکی از آن است که الگوریتم‌های مبتنی بر قواعد و درخت تصمیم، به‌ویژه M5Rules و Additive Regression، توانایی بالاتری در بازنمایی رفتار غیرخطی و نوسانات زمانی بارش دارند؛ موضوعی که در مقادیر بالاتر شاخص نش-ساتکلیف و خطای کمتر آن‌ها مشهود است.



شکل ۳- مقایسه روش Brute Force و میانگین ساده

Figure 3- Comparison of the Brute Force method and simple averaging

نتیجه‌گیری

نتایج این پژوهش نشان داد که داده‌های CHIRPS قادرند الگوی زمانی بارش ماهانه را در حوضه مورد مطالعه با دقت قابل‌قبولی بازنمایی کنند. مقادیر شاخص‌های ارزیابی حاکی از آن است که؛ این داده‌ها می‌توانند به‌عنوان منبع ورودی مناسب برای مدل‌های پیش‌بینی بارش مورد استفاده قرار گیرند. این موضوع نشان می‌دهد که؛ در شرایط محدودیت داده‌های ایستگاهی، استفاده از محصولات ماهواره‌ای می‌تواند نقش مؤثری در بهبود تحلیل‌های هیدرولوژیکی و توسعه مدل‌های پیش‌بینی ایفا کند. همچنین استفاده از رویکردهای ترکیب مدل‌ها می‌تواند دقت پیش‌بینی بارش ماهانه را نسبت به مدل‌های منفرد بهبود دهد. عملکرد بهتر مدل ترکیبی بیانگر آن است که؛ مدل‌های مختلف توانایی استخراج جنبه‌های متفاوتی از الگوهای زمانی بارش را دارند و ترکیب آن‌ها موجب کاهش خطاهای فردی مدل‌ها و افزایش پایداری پیش‌بینی می‌شود. این نتیجه با برخی مطالعات اخیر که به برتری روش‌های ensemble در مسائل پیش‌بینی بارش اشاره کرده‌اند هم‌راستا است (Ghosh et al., 2023; Lin et al., 2025). بنابراین می‌توان بیان کرد که بهره‌گیری هم‌زمان از داده‌های CHIRPS و رویکردهای ترکیب مدل‌ها، چارچوبی کارآمد برای بهبود پیش‌بینی بارش ماهانه در حوضه‌های دارای محدودیت داده‌های زمینی فراهم می‌کند.

ارزیابی دقت داده‌های CHIRPS با استفاده از چندین شاخص آماری شامل شاخص نش-ساتکلیف (NSE)، میانگین قدر مطلق خطا (MAE)، میانگین مربعات خطا (MSE) و شاخص بایاس (BIAS) نشان داد؛ که این پایگاه داده توانایی مناسبی در بازنمایی تغییرات زمانی بارش ماهانه دارد. ترکیب این شاخص‌ها بیانگر آن است که داده‌های CHIRPS علاوه بر داشتن همبستگی مناسب با داده‌های ایستگاهی، از نظر میزان خطا و سوگیری نیز عملکرد قابل‌قبولی دارند و می‌توانند به‌عنوان ورودی قابل اعتماد برای مدل‌های پیش‌بینی بارش مورد استفاده قرار گیرند. در مرحله مدل‌سازی، مقایسه عملکرد ۱۵ مدل مختلف یادگیری ماشین نشان داد که؛ توان پیش‌بینی مدل‌ها، تفاوت معناداری با یکدیگر دارد. در میان مدل‌های منفرد، الگوریتم‌های M5Rules، Additive Regression و Random Forest عملکرد بهتری نسبت به سایر مدل‌ها از خود نشان دادند. این برتری را می‌توان به توان بالای این مدل‌ها در شناسایی روابط غیرخطی، استخراج قواعد تصمیم‌گیری و کاهش خطاهای ساختاری نسبت داد. نتایج به‌دست‌آمده نشان می‌دهد که مدل‌های مبتنی بر درخت و قواعد، برای مدل‌سازی رفتار پیچیده و ناپایدار بارش ماهانه در مناطق خشک، کارایی بالاتری دارند.

همچنین یافته‌های اصلی پژوهش نشان داد که ترکیب مدل‌های

با این حال، نتایج نشان داد که حتی مدل‌هایی با عملکرد مناسب در حالت منفرد، در برخی بازه‌های زمانی دچار بایاس یا پراکندگی خطا می‌شوند. ترکیب دو به دوی مدل‌ها موجب کاهش این نارسایی‌ها شده و از طریق جبران متقابل خطاها، پایداری پیش‌بینی‌ها را افزایش داده است. توزیع مقادیر شاخص NSE برای ۱۰۵ ترکیب مختلف نشان داد که بخش قابل‌توجهی از ترکیب‌ها عملکردی بهتر از مدل‌های پایه داشته‌اند که بیانگر عمومی بودن مزیت رویکرد ترکیبی و عدم وابستگی آن به یک مدل خاص است.

مقایسه روش میانگین‌گیری ساده و وزن‌دهی BruteForce نشان داد که؛ اگرچه وزن‌دهی بهینه می‌تواند در برخی ترکیب‌ها منجر به بهبود بیشتر دقت شود، اما میانگین‌گیری ساده نیز بدون افزایش پیچیدگی محاسباتی، توانسته است به نتایجی رقابتی و پایدار دست یابد. این موضوع نشان می‌دهد که بهبود عملکرد صرفاً ناشی از افزایش پیچیدگی یا تعداد پارامترها نیست، بلکه از هم‌پوشانی ساختاری و مکمل بودن رفتار مدل‌ها حاصل می‌شود.

به‌طور کلی، برتری مدل‌های ترکیبی نسبت به مدل‌های منفرد را می‌توان با مبانی نظری تجمیع مدل‌ها توضیح داد. ترکیب چند مدل باعث می‌شود هر مدل سهمی از اطلاعات ورودی را که مدل‌های دیگر قادر به استخراج آن نیستند پوشش دهد و در نتیجه خطاهای ساختاری و بایاس مدل‌ها تا حدی یکدیگر را جبران کنند. افزون بر این، از آن‌جا که خطاهای مدل‌ها لزوماً همبستگی کامل ندارند، تلفیق خروجی‌ها منجر به کاهش واریانس پیش‌بینی و افزایش پایداری نتایج می‌گردد. در نهایت، وزن‌دهی بهینه در مدل ترکیبی موجب می‌شود سهم مدل‌هایی که عملکرد بهتری دارند افزایش یافته و مدل‌هایی با خطای بیشتر تأثیر کمتری در خروجی نهایی داشته باشند. مجموعه این عوامل موجب شده است که مدل ترکیبی در این مطالعه دقت بالاتری نسبت به مدل‌های منفرد ارائه دهد. یکی از دغدغه‌های اصلی در استفاده از مدل‌های ترکیبی، خطر بیش‌برازش است. در این پژوهش، برای کاهش این ریسک، چند اقدام اساسی انجام شد: نخست، تفکیک زمانی صریح داده‌ها به دوره آموزش و آزمون مستقل؛ دوم، ارزیابی عملکرد مدل‌ها صرفاً بر اساس داده‌های آزمون؛ و سوم، استفاده از روش میانگین‌گیری ساده که برخلاف ترکیب‌های پیچیده، احتمال بیش‌برازش را کاهش می‌دهد. همچنین پایداری نتایج در شاخص‌های مختلف و کاهش هم‌زمان MAE، MSE و بایاس نشان می‌دهد که بهبود عملکرد مدل‌ها ناشی از یادگیری الگوهای واقعی داده بوده. در مجموع، نتایج این بخش نشان می‌دهد که رویکرد ترکیبی پیشنهادی قادر است بدون ایجاد بیش‌برازش آشکار، دقت و پایداری پیش‌بینی بارش ماهانه را افزایش دهد.

وزن‌دهی مبتنی بر عملکرد مدل‌ها، میانگین وزنی پویا و سایر تکنیک‌های Ensemble، می‌تواند به بهبود دقت پیش‌بینی کمک کند. به کارگیری الگوریتم‌های پیشرفته‌تر یادگیری ماشین و یادگیری عمیق نیز می‌تواند در استخراج الگوهای زمانی پیچیده بارش مؤثر باشد. علاوه بر این، استفاده از متغیرهای ورودی بیشتر شامل پارامترهای هواشناسی و اقلیمی نظیر دما، رطوبت نسبی، فشار هوا و شاخص‌های اقلیمی بزرگ‌مقیاس، در کنار داده‌های سنجش‌ازدور، می‌تواند عملکرد مدل‌ها را بهبود دهد. ارزیابی روش‌های پیشنهادی در مناطق جغرافیایی و اقلیمی مختلف نیز برای بررسی میزان تعمیم‌پذیری مدل‌ها ضروری است. همچنین استفاده از روش‌های پیش‌پردازش داده، انتخاب ویژگی و کاهش ابعاد می‌تواند به کاهش نویز داده‌ها و افزایش دقت پیش‌بینی کمک کند. در نهایت، تحلیل عدم قطعیت مدل‌ها و بررسی حساسیت نتایج نسبت به متغیرهای ورودی می‌تواند به افزایش قابلیت اعتماد و کاربردپذیری نتایج در مطالعات منابع آب و مدیریت اقلیم کمک نماید.

مقایسه عملکرد مدل‌های مختلف نشان داد که؛ برخی از الگوریتم‌های یادگیری ماشین، به‌ویژه مدل‌های مبتنی بر قواعد مانند M5Rules و مدل‌های تقویتی نظیر Additive Regression، توانایی بیشتری در بازنمایی تغییرات زمانی بارش ماهانه دارند. همچنین نتایج نشان داد که استفاده از رویکرد ترکیب مدل‌ها موجب بهبود شاخص‌های ارزیابی نسبت به بسیاری از مدل‌های منفرد شده است. این موضوع بیانگر آن است که ترکیب مدل‌های مختلف می‌تواند با بهره‌گیری از قابلیت‌های مکمل آن‌ها، دقت پیش‌بینی را افزایش داده و پایداری نتایج را در مدل‌سازی بارش ماهانه بهبود بخشد.

یادگیری ماشین به‌صورت دو به دو، به‌طور کلی منجر به بهبود قابل توجه دقت پیش‌بینی نسبت به مدل‌های منفرد می‌شود. این بهبود ناشی از هم‌پوشانی نقاط قوت مدل‌ها و کاهش اثر خطاهای سیستماتیک هر یک از الگوریتم‌ها است. در میان تمامی ترکیب‌های بررسی‌شده، مدل ترکیبی Additive Regression - M5Rules با دستیابی به مقادیر شاخص نش-ساتکلیف ۰/۸۸، به‌عنوان بهترین ساختار پیش‌بینی شناسایی شد. کاهش قابل توجه خطا و بایاس پایین این مدل نشان‌دهنده پایداری و قابلیت اعتماد بالای آن در پیش‌بینی بارش ماهانه است. مقایسه این مدل با بهترین مدل‌های منفرد نیز نشان داد که حتی الگوریتم‌هایی با عملکرد مطلوب در حالت منفرد، در قالب ترکیبی می‌توانند دقت بالاتری ارائه دهند.

از منظر مدیریت منابع آب، نتایج این پژوهش اهمیت ویژه‌ای دارد. پیش‌بینی دقیق بارش ماهانه یکی از ورودی‌های اساسی در مدل‌های هیدرولوژیکی، مدیریت مخازن، برنامه‌ریزی تخصیص آب و ارزیابی خشکسالی است. بهبود دقت این پیش‌بینی‌ها می‌تواند منجر به کاهش عدم قطعیت در تصمیم‌گیری‌های مدیریتی، افزایش بهره‌وری بهره‌برداری از منابع آب و بهبود ارزیابی سناریوهای آینده اقلیمی شود. در حوضه‌هایی نظیر دریاچه نمک که با کمبود داده‌های زمینی مواجه‌اند، تلفیق داده‌های CHIRPS با مدل‌های ترکیبی یادگیری ماشین می‌تواند به‌عنوان یک چارچوب عملی و قابل تعمیم برای پشتیبانی از مدیریت پایدار منابع آب مورد استفاده قرار گیرد. با توجه به نتایج به‌دست‌آمده و محدودیت‌های موجود در این پژوهش، پیشنهاد می‌شود در مطالعات آینده از دوره‌های زمانی طولانی‌تر برای آموزش و به‌ویژه ارزیابی مدل‌ها استفاده شود تا پایداری و قابلیت تعمیم نتایج در شرایط اقلیمی مختلف بهتر بررسی گردد. همچنین بررسی و مقایسه روش‌های مختلف وزن‌دهی در مدل‌های ترکیبی، مانند روش‌های

References

1. Abdollahi, M., Abedi, K.J., & Matinzadeh, M. (2024). The effect of the sub-basins area and methods of calculating concentration time on the simulation of urban runoff volume using SewerGEMS software (Case study: Shahrekord). <https://doi.org/10.47176/jwss.28.3.49834>
2. Abinaya, P., & Janani, N. (2020). Rainfall forecasting using Weka data mining tool. *International Research Journal of Engineering and Technology*, 7(03).
3. Abishek, B., Priyatharshini, R., Eswar, M.A., & Deepika, P. (2017). Prediction of effective rainfall and crop water needs using data mining techniques. *2017 IEEE Technological Innovations in ICT for Agriculture and Rural Development (TIAR)*. <https://doi.org/10.1109/TIAR.2017.8273722>
4. Amrehn, M., Mualla, F., Angelopoulou, E., Steidl, S., & Maier, A. (2018). The random forest classifier in WEKA: Discussion and new developments for imbalanced data. *arXiv preprint arXiv:1812.08102*. <https://doi.org/10.48550/arXiv.1812.08102>
5. Beck, H.E., Pan, M., Roy, T., Weedon, G.P., Pappenberger, F., Van Dijk, A.I., Huffman, G.J., Adler, R.F., & Wood, E.F. (2019). Daily evaluation of 26 precipitation datasets using Stage-IV gauge-radar data for the CONUS. *Hydrology and Earth System Sciences*, 23(1), 207–224. <https://doi.org/10.5194/hess-23-207-2019>, 2019
6. Babaei Hesar, S., & Ghazavi, R. (2014). Comparison of TS and ANN models with the results of emission scenarios in rainfall prediction. *Journal of Soil and Water*. <https://doi.org/10.22067/jsw.v0i0.10261>
7. Biau, G., & Scornet, E. (2016). A random forest guided tour. *Test*, 25(2), 197–227. <https://doi.org/10.1007/s11749-016-0481-7>

8. Cervantes, J., Garcia-Lamont, F., Rodríguez-Mazahua, L., & Lopez, A. (2020). A comprehensive survey on support vector machine classification: Applications, challenges and trends. *Neurocomputing*, 408, 189–215. <https://doi.org/10.1016/j.neucom.2019.10.118>
9. Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 785–794). <https://doi.org/10.1145/2939672.2939785>
10. Dong, X., Yu, Z., Cao, W., Shi, Y., & Ma, Q. (2020). A survey on ensemble learning. *Frontiers of Computer Science*, 14(2), 241–258. <https://doi.org/10.1007/s11704-019-8208-z>
11. Fürnkranz, J., Gamberger, D., & Lavrač, N. (2012). *Foundations of Rule Learning*. Springer. <https://doi.org/10.1007/978-3-540-75197-7>
12. Fagan, M.E., & DeFries, R.S. (2024). Remote sensing and image processing. In S. M. Scheiner (Ed.), *Encyclopedia of Biodiversity* (3rd ed., pp. 432–445). Academic Press. <https://doi.org/10.1016/B978-0-12-822562-2.00060-8>
13. Fallah-Ghalhari, Gh.A., Mousavi Baygi, M., & Habibi Nokhandan, M. (2009). Results compression of Mamdani fuzzy interface system and artificial neural networks in the seasonal rainfall prediction (Case study: Khorasan region).
14. Funk, C., Peterson, P., Landsfeld, M., Pedreros, D., Verdin, J., Shukla, S., Husak, G., Rowland, J., Harrison, L., & Hoell, A. (2015). The climate hazards infrared precipitation with stations—a new environmental record for monitoring extremes. *Scientific Data*, 2(1), 1–21. <https://doi.org/10.1038/sdata.2015.66>
15. Fikileni, S., & Wolski, P. (2022). Framework for implementation of the Pitman-WR2012 model in seasonal hydrological forecasting: A case study of Kraai River, South Africa. *Water SA*, 48(1), 62–74. <https://doi.org/10.17159/WSA/2022.V48.I1.3891>
16. Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep Learning*. MIT Press. <https://doi.org/10.4258/hir.2016.22.4.351>
17. Ghosh, S., Gourisaria, M.K., Sahoo, B., & Das, H. (2023). A pragmatic ensemble learning approach for rainfall prediction. *Discover Internet of Things*, 3(13). <https://doi.org/10.1007/s43926-023-00044-3>
18. Haddadi, A., & Yazdi, M. (2025). Performance assessment of the extreme gradient boosting machine model in forecasting precipitation depth to improve precipitation estimation accuracy in data-scarce regions. *Journal of Irrigation and Water Engineering*. <https://doi.org/10.22034/iwrr.2025.543507.2944>
19. Holmes, G., Hall, M., & Frank, E. (1999). Generating rule sets from model trees. *Australian Joint Conference on Artificial Intelligence* (pp. 1–12). Springer. https://doi.org/10.1007/3-540-46695-9_1
20. James, G., Witten, D., Hastie, T., & Tibshirani, R. (2021). *An Introduction to Statistical Learning: with Applications in R* (2nd ed.). Springer. <https://doi.org/10.1007/978-1-4614-7138-7>
21. Kumar, P.S., & Swathi, M.N. (2023). Rainfall prediction using ada boost machine learning ensemble algorithm. *Journal of Advanced Applied Scientific Research*, 5(4), 67–81. <https://doi.org/10.46947/joaasr542023682>
22. Kuncheva, L.I. (2014). *Combining Pattern Classifiers: Methods and Algorithms* (2nd ed.). John Wiley & Sons. <https://doi.org/10.1109/TNN.2007.897478>
23. Mahjoobi, E., & Rafiei, F. (2025). Exploring the potential of geospatial data: An in-depth. In *Exploring Remote Sensing-Methods and Applications* (p. 113). <https://doi.org/10.5772/intechopen.1006999>
24. Lin, S.-S., Zhu, K.-Y., Wang, C.-Y., Yang, C.-P., & Liu, M.-Y. (2025). Ensemble Learning-Based Soft Computing Approach for Future Precipitation Analysis. *Atmosphere*, 16(6), 669. <https://doi.org/10.3390/atmos16060669>
25. Oswal, N. (2019). Predicting rainfall using machine learning techniques. *arXiv preprint arXiv:1910.13827*. <https://doi.org/10.48550/arXiv.1910.13827>
26. Omary, A., Wedyan, A., Zghoul, A., Banihani, A., & Alsmadi, I. (2012). An interactive predictive system for weather forecasting. *2012 International Conference on Computer, Information and Telecommunication Systems (CITS)*. <https://doi.org/10.1109/CITS.2012.6220375>
27. Patel, A., Keriwala, N., Soni, N., Goel, U., Bhoj, R., Adhyaru, Y., & Yadav, S. (2023). Rainfall prediction using machine learning techniques for Sabarmati River Basin, Gujarat, India. *Journal of Engineering Science & Technology Review*, 16(1). <https://doi.org/10.25103/jestr.161.13>
28. Ramani, K., Reddy, M.S., Bhavani, K., Feeza, S., & Baves, V.S. (2024). Optimization of rainfall prediction using satellite data through machine learning and deep learning algorithms. *2024 IEEE International Conference on Information Technology, Electronics and Intelligent Communication Systems (ICITEICS)*. <https://doi.org/10.1109/ICITEICS61368.2024.10625624>
29. Rezaei, M., Nohtani, M., Moghaddamnia, A., Abkar, A., & Rezaei, M. (2014). Performance evaluation of Statistical Downscaling Model (SDSM) in forecasting precipitation in two arid and hyper arid regions. *Journal of Soil and Water*. <https://doi.org/10.22067/jsw.v0i0.23119>
30. Rahman, A.-u., Abbas, S., Gollapalli, M., Ahmed, R., Aftab, S., Ahmad, M., Khan, M.A., & Mosavi, A. (2022). Rainfall prediction system using machine learning fusion for smart cities. *Sensors*, 22(9), 3504. <https://doi.org/10.3390/s22093504>

31. Schulz, E., Speekenbrink, M., & Krause, A. (2018). A tutorial on Gaussian process regression: Modelling, exploring, and exploiting functions. *Journal of Mathematical Psychology*, 85, 1–16. <https://doi.org/10.1016/j.jmp.2018.03.001>
32. Salahi, A., Ashrafzadeh, A., & Vazifedoust, M. (2024a). Assessing the forecasting accuracy of intense precipitation events in Iran using the WRF model. *Earth Sciences Informatics*, 17, 2199–2211. <https://doi.org/10.1007/s12145-024-01274-x>
33. Salahi, A., Ashrafzadeh, A., & Vazifedoust, M. (2024b). Remote sensing-based precipitation forecasting using cloud optical characteristics: threshold optimization and evaluation in Northern and Western Iran. *Natural Hazards*, 120, 3661–3675. <https://doi.org/10.1007/s11069-023-06352-9>
34. Sagi, O., & Rokach, L. (2018). Ensemble learning: A survey. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 8(4), e1249. <https://doi.org/10.1002/widm.1249>
35. Tuysuzoglu, G., Birant, K.U., & Birant, D. (2023). Rainfall prediction using an ensemble machine learning model based on K-stars. *Sustainability*, 15(7), 5889. <https://doi.org/10.3390/su15075889>
36. Taunk, K., De, S., Verma, S., & Swetapadma, A. (2019). A brief review of nearest neighbor algorithm for learning and classification. In *2019 International Conference on Intelligent Computing and Control Systems (ICCS)* (pp. 1255–1260). IEEE. <https://doi.org/10.1109/ICCS45141.2019.9065747>
37. Tukey, J.W. (1977). Exploratory Data Analysis. https://doi.org/10.1007/978-3-031-20719-8_2
38. Velasco, L.C., Aca-ac, J.M., Cajés, J.J., Lactuan, N.J., & Chit, S.C. (2022). Rainfall forecasting using support vector regression machines. *International Journal of Advanced Computer Science and Applications*, 13(3). <https://doi.org/10.14569/IJACSA.2022.0130329>
39. Wang, Y., & Witten, I.H. (1997). Induction of model trees for predicting continuous classes. *Proceedings of the Poster Papers of the European Conference on Machine Learning* (pp. 48–53). <https://hdl.handle.net/10289/1183>
40. Wu, Y., Wang, H., Zhang, B., & Du, K.L. (2012). Using radial basis function networks for function approximation and classification. *ISRN Applied Mathematics*. <https://doi.org/10.5402/2012/324194>
41. Xu, J., & Zheng, Y. (2024). Remote sensing data analysis for urban planning and land use change. *Trans. Environment Energy Earth Science*, 3, 20–25. <https://doi.org/10.62051/gwpgk5m73>