



ORIGINAL RESEARCH PAPER

Yield prediction of winter wheat using remote sensing data and machine-learning algorithms

F. Babayi, A. Vafaei-Nejad*, A. Sharifi

Department of Surveying Engineering, Faculty of Civil, Water and Environmental Engineering, Shahid Beheshti University, Tehran, Iran

ABSTRACT

Received: 24 February 2026

Reviewed: 17 April 2026

Revised: 04 June 2026

Accepted: 11 August 2026

KEYWORDS:

Winter Wheat Yield

Remote Sensing

XGBoost

SHAP

NDVI

VPD

* Corresponding author

✉ a_vafaei@sbu.ac.ir

☎ (+9821) 73932452

Background and Objectives: With the growing global population, increasing pressure on natural resources, climate change, and the rising importance of food security, accurate crop yield prediction—particularly for wheat—has become a crucial research topic. In recent years, the integration of remote sensing data, climatic information, and soil properties with machine-learning algorithms has created new opportunities for estimating crop yields. However, many previous studies have either focused on a single data source or have been conducted at limited spatial and temporal scales. Consequently, a gap still exists in conducting comprehensive and comparative evaluations of models using multi-source datasets. The aim of this study was to evaluate the capability of different machine-learning algorithms to predict wheat yield using climatic variables and the remotely sensed Normalized Difference Vegetation Index (NDVI), and to assess the generalizability of these models across different spatial and climatic conditions.

Methods: This study was conducted using multi-source data comprising climatic variables and the remotely sensed Normalized Difference Vegetation Index (NDVI). The study covered nine regions in Germany over a 22-year period, thereby enabling the assessment of spatial and temporal variations in wheat yield. The XGBoost, Random Forest, and Ridge Regression algorithms were employed to predict wheat yield. Among the investigated variables, NDVI and vapor pressure deficit (VPD), as an indicator of atmospheric evaporative stress, were identified as important predictors in the model analyses. Model performance was evaluated in terms of predictive accuracy and error using the coefficient of determination (R^2), root mean square error (RMSE), and mean absolute error (MAE). The study was designed not only to analyze the data but also to enable a regional comparison between conventional and advanced predictive models.

Findings: The results showed that the XGBoost algorithm outperformed both Random Forest and Ridge Regression in predicting wheat yield. XGBoost was able to reconstruct spatial and temporal yield patterns with lower error and higher explanatory power. The findings also indicated that NDVI and vapor pressure deficit (VPD), as an indicator of atmospheric evaporative stress, were important predictors of wheat yield patterns. Higher model accuracy was observed in more homogeneous and high-yielding regions, whereas prediction errors increased in areas with greater environmental heterogeneity. Nevertheless, the overall trend demonstrated that machine-learning techniques—especially gradient-boosting models—have strong capabilities in capturing the complex relationships between environmental factors and crop yield.

Conclusion: Overall, the findings of this study demonstrate that integrating remote-sensing and climatic data with advanced machine-learning algorithms can provide an effective tool for predicting wheat yield and supporting decision-making in smart agriculture. The superior performance of XGBoost compared with the other models indicates its strong capability to capture the nonlinear and multidimensional relationships inherent in agricultural data. Nevertheless, this study had several limitations, including the lack of farm-management data, such as sowing dates, cultivar types, irrigation levels, and fertilizer application rates, as well as the absence of soil data with sufficient spatial and temporal resolution for inclusion in the final models. Moreover, the use of more complex models may increase the risk of overfitting when data are limited. Therefore, future studies are recommended to employ more comprehensive datasets, hybrid spatiotemporal models, and advanced approaches such as recurrent neural networks and convolutional neural networks. Such approaches could improve predictive accuracy and model generalizability while enhancing the practical application of yield-prediction models in sustainable agricultural management and food-security planning.



NUMBER OF REFERENCES

37



NUMBER OF FIGURES

15



NUMBER OF TABLES

4

مقاله پژوهشی

پیش‌بینی عملکرد محصول گندم زمستانه با استفاده از داده‌های سنجش‌ازدور و الگوریتم‌های یادگیری ماشین

فاطمه بابایی، علیرضا وفائی نژاد*، علیرضا شریفی

گروه مهندسی نقشه برداری، دانشکده مهندسی عمران آب و محیط‌زیست، دانشگاه شهید بهشتی، تهران، ایران

چکیده

پیشینه و اهداف: باتوجه‌به افزایش جمعیت، فشار بر منابع طبیعی، تغییرات اقلیمی و اهمیت روزافزون امنیت غذایی، پیش‌بینی دقیق عملکرد محصولات کشاورزی، به‌ویژه گندم، به یکی از موضوعات مهم پژوهشی تبدیل شده است. در سال‌های اخیر، استفاده از داده‌های سنجش‌ازدور، اطلاعات اقلیمی و ویژگی‌های خاک در کنار الگوریتم‌های یادگیری ماشین، فرصت‌های جدیدی برای برآورد عملکرد محصول فراهم کرده است. بااین‌حال، بسیاری از پژوهش‌های پیشین یا تنها بر یک منبع داده تمرکز داشته‌اند یا در مقیاس‌های محدود مکانی و زمانی انجام شده‌اند؛ بنابراین هنوز شکافی در زمینه ارزیابی جامع و مقایسه‌ای مدل‌ها با استفاده از داده‌های چندمنبعی وجود دارد. هدف این پژوهش، بررسی توانایی الگوریتم‌های مختلف یادگیری ماشین در پیش‌بینی عملکرد گندم با بهره‌گیری از داده‌های اقلیمی، شاخص NDVI حاصل از سنجش‌ازدور و همچنین ارزیابی قابلیت تعمیم این مدل‌ها در شرایط مکانی و اقلیمی متفاوت بوده است.

روش‌ها: این پژوهش بر اساس داده‌های چندمنبعی شامل متغیرهای اقلیمی و شاخص سنجش‌ازدور NDVI انجام شد. مطالعه در ۹ منطقه مختلف در آلمان و در بازه‌ای ۲۲ ساله صورت‌گرفته و بدین ترتیب امکان بررسی تغییرات مکانی و زمانی عملکرد گندم فراهم شده است. در این تحقیق از الگوریتم‌های XGBoost، Random Forest و Ridge Regression برای پیش‌بینی عملکرد گندم استفاده شده است. از میان متغیرهای مؤثر، شاخص NDVI و تنش تبخیری از جمله عوامل مهم در تحلیل مدل‌ها بوده‌اند. عملکرد مدل‌ها از نظر دقت پیش‌بینی و میزان خطا مورد ارزیابی قرار گرفته و شاخص‌هایی مانند ضریب تعیین، ریشه میانگین مربعات خطا و میانگین خطای مطلق برای مقایسه نتایج استفاده شده‌اند. طراحی مطالعه به‌گونه‌ای بوده که علاوه بر تحلیل داده‌ها، امکان مقایسه عملکرد مدل‌های کلاسیک و پیشرفته در سطح منطقه‌ای نیز فراهم شود.

یافته‌ها: نتایج پژوهش نشان داد که الگوریتم XGBoost نسبت به دو مدل Ridge Regression و Random Forest عملکرد بهتری در پیش‌بینی عملکرد گندم داشته است. این مدل توانست الگوهای فضایی و زمانی عملکرد را با خطای کمتر و قدرت تبیین بالاتر بازسازی کند. همچنین مشخص شد که متغیرهای NDVI، تنش تبخیری نقش مهمی در شکل‌گیری الگوهای عملکرد دارند. در مناطق یکنواخت‌تر و پربازده‌تر، دقت مدل بالاتر بود، درحالی‌که در مناطقی با ناهمگنی محیطی بیشتر، مانند برخی نواحی خاص، خطای پیش‌بینی افزایش یافت. باوجوداین، روند کلی نشان داد که روش‌های یادگیری ماشین، به‌ویژه مدل‌های مبتنی بر تقویت‌گرادیان، توان بالایی در شناسایی روابط پیچیده بین عوامل محیطی و عملکرد محصول دارند.

نتیجه‌گیری: در مجموع، نتایج این پژوهش نشان می‌دهد که ترکیب داده‌های سنجش‌ازدور و اقلیمی همراه با الگوریتم‌های پیشرفته یادگیری ماشین می‌تواند ابزاری مؤثر برای پیش‌بینی عملکرد گندم و پشتیبانی از تصمیم‌گیری در کشاورزی هوشمند باشد. برتری XGBoost نسبت به سایر مدل‌ها حاکی از آن است که مدل‌های مبتنی بر یادگیری ماشین توانایی بالایی در تحلیل روابط غیرخطی و چندبعدی موجود در داده‌های کشاورزی دارند. بااین‌حال، محدودیت‌هایی نیز وجود دارد؛ از جمله نبود داده‌های مدیریتی مزرعه مانند تاریخ کاشت، نوع رقم، میزان آبیاری و مصرف کود، و نبود داده‌های خاکی با تفکیک مکانی و زمانی کافی برای ورود به مدل نهایی، از محدودیت‌های پژوهش حاضر بودند. همچنین استفاده از مدل‌های پیچیده‌تر ممکن است در شرایط داده محدود با خطر بیش‌برازش همراه باشد؛ بنابراین، پیشنهاد می‌شود در پژوهش‌های آینده از مجموعه داده‌های کامل‌تر، مدل‌های هیبریدی فضایی - زمانی و روش‌های پیشرفته‌تر مانند شبکه‌های عصبی بازگشتی و کانولوشنی استفاده شود. این رویکرد می‌تواند به افزایش دقت پیش‌بینی، تعمیم‌پذیری بیشتر و کاربرد عملی بهتر در مدیریت پایدار کشاورزی و امنیت غذایی کمک کند.

تاریخ دریافت: ۰۵ اسفند ۱۴۰۴
تاریخ داوری: ۲۸ فروردین ۱۴۰۵
تاریخ اصلاح: ۱۴ خرداد ۱۴۰۵
تاریخ پذیرش: ۲۰ مرداد ۱۴۰۵

واژگان کلیدی:

عملکرد گندم زمستانه
سنجش‌ازدور
XGBoost
SHAP
NDVI
VPD

* نویسنده مسئول

a_vafaei@sbu.ac.ir

۰۲۱-۷۳۹۳۲۴۵۲ ①

مقدمه

امنیت غذایی یکی از چالش‌های کلیدی قرن بیست و یکم است که تحت‌تأثیر تغییرات اقلیمی، رشد جمعیت و کاهش منابع طبیعی قرار دارد. غلات به‌عنوان منابع اصلی تغذیه نقش حیاتی دارند. افزایش تقاضا و محدودیت زمین‌های کشاورزی ضرورت تولید پایدار را دوچندان کرده است. در دهه‌های اخیر، مدل‌سازی عملکرد محصول بر پایه داده‌های اقلیمی مانند دما، بارش، تابش خورشیدی و فشار بخار اشباع موردتوجه قرار گرفته است. مدل‌های آماری کلاسیک محدودیت‌هایی در درک روابط غیرخطی دارند، درحالی‌که الگوریتم‌های یادگیری ماشین و یادگیری عمیق قادر به کشف الگوهای پنهان هستند و ترکیب این دو رویکرد می‌تواند دقت و تعمیم‌پذیری مدل‌ها را افزایش دهد.

باتوجه به اهمیت روزافزون پیش‌بینی عملکرد محصولات کشاورزی و نقش فناوری‌های داده‌محور در بهبود دقت این پیش‌بینی‌ها، در سال‌های اخیر پژوهش‌های متعددی در زمینه توسعه مدل‌های یادگیری ماشین و یادگیری عمیق انجام شده است. در ادامه، ابتدا زمینه و ضرورت پژوهش تشریح شده و سپس مرور مطالعات پیشین و خلأهای پژوهشی ارائه می‌شود.

امنیت غذایی یکی از چالش‌های کلیدی قرن بیست و یکم است که تحت‌تأثیر رشد جمعیت، تغییرات اقلیمی، کاهش منابع آب‌و‌خاک و افزایش تقاضا قرار دارد. پیش‌بینی دقیق عملکرد محصولات، نقش مهمی در مدیریت مزرعه، سیاست‌گذاری، بازار و کاهش ریسک‌های اقلیمی ایفا می‌کند. در گذشته، برآورد عملکرد عمدتاً با مدل‌های آماری کلاسیک یا مدل‌های رشد گیاه انجام می‌شد؛ اما این روش‌ها به دلیل فرض خطی بودن و ناتوانی در مدل‌سازی تعاملات پیچیده، دقت کافی نداشتند. توسعه سنجش‌از‌دور، دسترسی به داده‌های اقلیمی با تفکیک مناسب و پیشرفت یادگیری ماشین، رویکردهای نوینی فراهم کرده است [۱، ۲]. تلفیق داده‌های سنجش‌از‌دور و اقلیمی، یکی از مؤثرترین راهکارها برای افزایش دقت پیش‌بینی است. داده‌های اقلیمی شرایط محیطی و داده‌های سنجش‌از‌دور وضعیت فنولوژیکی، سلامت و تغییرات بیوفیزیکی گیاه را نشان می‌دهند. استفاده هم‌زمان این دو منبع، دقت، پایداری و قابلیت تعمیم مدل‌ها را افزایش می‌دهد [۳، ۴]. تصاویر Sentinel-2 با قدرت تفکیک مناسب، استخراج شاخص‌هایی نظیر NDVI، EVI و GNDVI را ممکن ساخته‌اند که تغییرات رشد، سبزیگی، فتوسنتز و تنش‌ها را پایش می‌کنند و در ترکیب با داده‌های اقلیمی، دقت پیش‌بینی را بهبود می‌بخشند [۵، ۶]. داده‌های چندمنبعی اطلاعات مکمل از جنبه‌های مختلف رشد گیاه ارائه می‌دهند و ترکیب آن‌ها در مدل‌های داده‌محور، علاوه بر دقت، توان تعمیم را افزایش می‌دهد [۳، ۷].

توسعه الگوریتم‌های یادگیری ماشین مانند Random Forest، SVM، Gradient Boosting و XGBoost، به دلیل توانایی در مدل‌سازی غیرخطی و مدیریت داده‌های پیچیده، دقت بیشتری نسبت به روش‌های کلاسیک دارند [۲، ۸، ۹]. مرور نظام‌مند نشان می‌دهد که دما، بارندگی،

ویژگی‌های خاک و شاخص‌های پوشش گیاهی مهم‌ترین ورودی‌ها و RMSE، MAE و R^2 رایج‌ترین معیارهای ارزیابی هستند [۱۰]. پژوهش‌های مروری نیز حاکی از افزایش سهم یادگیری عمیق و تلفیق داده‌های چندمنبعی است [۱۱]. مدل‌های CNN، LSTM و مبتنی بر توجه، الگوهای مکانی-زمانی پیچیده را استخراج می‌کنند. برای نمونه، Jorvekar و همکاران (۲۰۲۴) با CNN و تلفیق داده‌های حسگرها و سوابق عملکرد، دقت بالایی گزارش کردند. El Ayyadi و همکاران (۲۰۲۶) با چارچوب STAFNet و ترکیب Sentinel-2، اقلیم و توجه، به R^2 حدود ۰٫۹۳ دست یافتند و نشان دادند داده‌های مصنوعی GAN تعمیم را بهبود می‌بخشد. مدل SCM-GAT نیز با وارد کردن روابط علی، دقت و تفسیرپذیری را افزایش داد [۱۲].

با وجود پیشرفت‌ها، استفاده عملی از یادگیری عمیق با چالش‌هایی مانند نیاز به داده‌های زیاد، بیش‌برازش در مجموعه‌های کوچک و کاهش تفسیرپذیری همراه است؛ بنابراین انتخاب مدل باید متناسب با داده باشد [۱۳، ۱۴]. از این‌رو، الگوریتم‌های مبتنی بر درخت تصمیم مانند Random Forest و XGBoost همچنان کارآمد هستند. هم‌زمان، تفسیرپذیری با روش‌هایی مانند SHAP و Permutation Importance به یکی از محورهای اصلی تبدیل شده است که امکان شناسایی متغیرهای اثرگذار و افزایش اعتمادپذیری را فراهم می‌کند [۱۵، ۱۶]. Srivastava و همکاران (۲۰۲۲) در مطالعه گندم آلمان، نشان دادند CNN دقیق‌ترین است اما SHAP متغیرهای کلیدی را شناسایی می‌کند و همچنین در ایران، جهان (۲۰۲۵) با SHAP نشان داد دمای کمینه و بارندگی مهم‌ترین عوامل در Random Forest و XGBoost هستند [۱۶، ۱۷].

مطالعات کاربردی در مراکش [۴]، سوئد [۱۸] و لیتوانی [۵] بر اهمیت داده‌های چندمنبعی تأکید دارند. در آلمان، Sahu و همکاران (۲۰۲۶) نشان دادند تعمیم مکانی مدل‌ها نیازمند بازآموزی منطقه‌ای است [۱۹] و محمدی و همکاران (۲۰۲۳) بر مدیریت منطقه‌محور تأکید کردند [۲۰]. در ایران، مطالعات متعدد [۲۱-۲۸] نشان داده‌اند که مدل‌های غیرخطی و ترکیبی، به‌ویژه با شاخص‌های طیفی، دقت پیش‌بینی را بهبود می‌بخشند. با این حال، بیشتر این پژوهش‌ها در مقیاس محلی یا استانی، با یک منبع داده یا تعداد محدودی الگوریتم انجام شده‌اند و ارزیابی هم‌زمان مدل‌های رگرسیونی، Ensemble و تفسیرپذیری با داده‌های چندمنبعی و بلندمدت کمتر مورد توجه بوده است.

مرور مطالعات نشان می‌دهد که پیش‌بینی عملکرد از مدل‌های سنتی به سمت یادگیری ماشین و عمیق حرکت کرده و تلفیق داده‌های اقلیمی، سنجش‌از‌دور و محیطی دقت را افزایش می‌دهد [۳، ۷، ۹، ۱۱]. روش‌های Explainable AI نیز تفسیرپذیری را ممکن ساخته‌اند [۱۷-۱۵]. با وجود این، خلأهایی شامل تمرکز بر یک منبع داده، مقیاس محدود مکانی-زمانی، کمبود ارزیابی تعمیم‌پذیری و چالش‌های یادگیری عمیق در کاربرد عملی، همچنان وجود دارد [۳، ۱۳، ۱۴، ۱۸-۲۰]. در ایران نیز بیشتر پژوهش‌ها در مقیاس محدود و با تعداد الگوریتم کم بوده

با وجود اینکه مدل‌های یادگیری عمیق مانند LSTM و CNN نیز در پیش‌بینی عملکرد محصولات کشاورزی نتایج قابل‌قبولی ارائه کرده‌اند، این مدل‌ها معمولاً برای دستیابی به عملکرد مطلوب به حجم زیادی از داده‌های آموزشی و یا داده‌های دارای ساختار زمانی یا تصویری نیاز دارند. در پژوهش حاضر، مجموعه داده شامل ۲۳۳۲ نمونه ساخت‌یافته از ۹ منطقه آلمان طی یک دوره ۲۲ ساله بود که برای آموزش مدل‌های عمیق نسبتاً محدود محسوب می‌شود و می‌تواند خطر بیش‌برازش را افزایش دهد. از سوی دیگر، داده‌های مورد استفاده عمدتاً به صورت جدولی بودند و شامل تصاویر خام ماهواره‌ای یا سری‌های زمانی پیوسته مورد نیاز مدل‌های CNN و LSTM نبودند؛ بنابراین، استفاده از الگوریتم‌های مبتنی بر درخت تصمیم، به‌ویژه XGBoost و Random Forest، انتخاب مناسب‌تری برای این نوع داده‌ها محسوب می‌شود؛ زیرا این الگوریتم‌ها ضمن ارائه دقت بالا، نیاز کمتری به حجم زیاد داده داشته و از قابلیت تفسیرپذیری بیشتری نیز برخوردار هستند.

مواد و روش‌ها

منطقه مورد مطالعه

آلمان در اروپای مرکزی واقع شده و تحت تأثیر ترکیبی از اقلیم اقیانوسی در بخش‌های شمالی و غربی و اقلیم قاره‌ای در بخش‌های شرقی و جنوبی قرار دارد. این تنوع اقلیمی، همراه با تفاوت‌های ارتفاعی و ویژگی‌های خاک، باعث می‌شود شرایط رشد گندم در مناطق مختلف کشور متفاوت باشد؛ به‌طوری که الگوهای دما، بارش، تابش و تنش‌های آبی - حرارتی بین شمال مرطوب و نسبتاً خنک، تا جنوب و جنوب شرق با شرایط نسبتاً گرم‌تر و گاه خشک‌تر تغییر می‌کند.

در این پژوهش، ۹ منطقه مختلف در آلمان به عنوان مناطق مطالعه انتخاب شده‌اند که هر یک نماینده بخشی از این گرادیان اقلیمی و محیطی هستند. این مناطق در بخش‌های شمالی (Hamburg، Schleswig، Rostock)، مرکزی (Kassel، Erfurt-Weimar، Meiningen) و جنوبی/جنوب شرقی (Augsburg، Weiden، Freudenstadt) آلمان قرار دارند و به این ترتیب، طیفی از شرایط اقلیمی دریایی تا قاره‌ای و از ارتفاعات پایین تا مناطق مرتفع‌تر را پوشش می‌دهند.

مجموعه داده مورد استفاده در این مقاله شامل ۲۳۳۲ نمونه از این ۹ منطقه است که مربوط به ۲۲ سال مختلف است. توزیع داده‌ها بین مناطق به گونه‌ای است که Hamburg و Schleswig هر یک ۱۳٫۲۱٪ از کل داده‌ها را به خود اختصاص داده‌اند؛ مناطق Erfurt-Weimar، Rostock و Weiden هر کدام ۱۱٫۳۲٪ از نمونه‌ها را شامل می‌شوند؛ Augsburg، Freudenstadt، Meiningen هر یک ۱۰٫۳۸٪ و Kassel، ۸٫۴۹٪ از داده‌ها را در بر دارد. این توزیع، پوشش نسبتاً متعادلی از شرایط اقلیمی و محیطی مختلف در آلمان فراهم می‌کند و امکان تحلیل مکانی روابط بین متغیرهای اقلیمی و عملکرد گندم در چارچوب خروجی مدل سراسری را مهیا می‌سازد.

و تحلیل تفسیرپذیری و تعمیم‌پذیری کمتر مدنظر قرار گرفته است [۲۴، ۲۶، ۲۷].

گرچه پژوهش‌های متعددی در زمینه پیش‌بینی عملکرد محصولات کشاورزی با استفاده از روش‌های یادگیری ماشین و یادگیری عمیق انجام شده است، بخش قابل‌توجهی از پژوهش‌های پیش‌بینی عملکرد محصولات بر مناطق خاصی مانند آفریقای سیاه، جنوب آسیا یا اقلیم‌های خشک و نیمه‌خشک متمرکز بوده‌اند و کاربرد چارچوب‌های پیشنهادی در اقلیم‌های معتدل اروپایی، از جمله آلمان، کمتر مورد بررسی قرار گرفته است. این مسئله، تعمیم‌پذیری مکانی و اقلیمی نتایج را محدود می‌کند و مانعی برای ارزیابی رفتار مدل‌ها در شرایط واقعی مزارع تحت اقلیم‌های معتدل به‌شمار می‌آید [۴، ۲۹، ۳۰].

بسیاری از مطالعات یا صرفاً از مدل‌های آماری کلاسیک (مانند رگرسیون خطی چندگانه و رگرسیون چندجمله‌ای) استفاده کرده‌اند یا تنها بر الگوریتم‌های یادگیری ماشین (مانند Random Forest، XGBoost و سایر روش‌های تجمیعی) تمرکز داشته‌اند. در عین حال، مقایسه‌ای نظام‌مند و یکپارچه میان این دو دسته روش، در یک چارچوب مشترک که داده‌های اقلیمی واقعی و شاخص‌های سنجش‌ازدور چندزمانه را هم‌زمان در برگیرد، کمتر گزارش شده است. این خلأ باعث می‌شود ارزیابی دقیق مزیت نسبی مدل‌های داده‌محور نسبت به روش‌های کلاسیک، به‌ویژه در اقلیم‌های معتدل، دشوار باشد [۳۱-۳۴]. همین راستا، پژوهش حاضر با هدف پر کردن این خلأ طراحی شده است. در این مطالعه، مجموعه‌ای گسترده از داده‌های اقلیمی، عملکرد نهایی دانه و شاخص‌های سنجش‌ازدور برای ۹ منطقه مختلف در کشور آلمان طی یک دوره زمانی ۲۲ ساله مورد استفاده قرار گرفته است. این گستره زمانی و مکانی امکان انجام یک تحلیل سراسری و مقایسه‌پذیر را فراهم می‌کند و شرایط متنوع اقلیمی و مدیریتی را در بر می‌گیرد. نتایج این مطالعه می‌تواند درک دقیق‌تری از قابلیت تعمیم‌پذیری مدل‌های پیش‌بینی عملکرد در شرایط متنوع اقلیمی فراهم کرده و بینش‌های کاربردی برای توسعه کشاورزی هوشمند، مدیریت بهینه منابع و تقویت امنیت غذایی ارائه دهد.

در این پژوهش، سه مدل Ridge Regression، Random Forest و XGBoost با هدف مقایسه سه رویکرد متفاوت در مدل‌سازی داده‌های ساخت‌یافته انتخاب شدند. مدل Ridge Regression به عنوان یک روش رگرسیونی خطی همراه با منظم‌سازی (Regularization)، مبنایی برای ارزیابی عملکرد مدل‌های پیچیده‌تر فراهم می‌کند. در مقابل، Random Forest به عنوان یکی از شناخته‌شده‌ترین الگوریتم‌های مبتنی بر Bagging، توانایی مناسبی در مدل‌سازی روابط غیرخطی و کاهش واریانس پیش‌بینی دارد. همچنین XGBoost به عنوان یکی از پیشرفته‌ترین الگوریتم‌های Boosting، به دلیل دقت بالا، توانایی مدیریت روابط پیچیده بین متغیرها و مقاومت مناسب در برابر بیش‌برازش، در سال‌های اخیر به طور گسترده در مطالعات پیش‌بینی عملکرد محصولات کشاورزی مورد استفاده قرار گرفته است.



شکل ۱: موقعیت مناطق مورد مطالعه در آلمان برای پیش‌بینی عملکرد گندم زمستانه.

Fig. 1: Location of the study areas in Germany for winter wheat yield prediction.

داده‌ها

داده‌های اقلیمی مورد استفاده در این پژوهش از مجموعه داده ارائه شده در مطالعه Schils و همکاران (۲۰۱۸) درباره شکاف عملکرد غلات در اروپا استخراج شد. در این مطالعه، داده‌های مربوط به یکی از مناطق مطالعاتی آن پژوهش مورد استفاده قرار گرفت و برای تکمیل مجموعه داده، شاخص سنجش از دور NDVI نیز به متغیرهای اقلیمی افزوده شد [۳۵].

داده‌های عملکرد گندم

متغیر هدف در این پژوهش، عملکرد نهایی دانه گندم (Yield)، بر حسب kg/ha است که برای هر رکورد منطقه - سال در مجموعه داده ثبت شده است. این مجموعه شامل ۲۳۳۲ نمونه از ۹ منطقه آلمان است. مقدار عملکرد دامنه‌ای بین ۶۲۰۶ تا ۱۰۶۰۱ kg/ha داشته و میانگین آن ۸۵۷۶٫۶ kg/ha برآورد شده است. ضریب تغییرات حدود ۱۱٫۸٪ نشان می‌دهد که اگرچه عملکرد در سطح کشور نسبتاً یکنواخت است، اما

همچنان نوسانات مهمی ناشی از تفاوت‌های اقلیمی و محیطی وجود دارد. این متغیر به عنوان خروجی اصلی در تمامی مدل‌های رگرسیونی و یادگیری ماشین استفاده شده است.

داده‌های اقلیمی

برای توصیف شرایط حرارتی - آبی مؤثر بر رشد گندم، از شش متغیر اقلیمی اصلی استفاده شد:

Tmax: میانگین سالانه دمای بیشینه (°C)

Tmin: میانگین سالانه دمای کمینه (°C)

VPD: کسری فشار بخار (kPa)

SR: مجموع تابش خورشیدی (MJ m^{-2})

prcp: مجموع بارش سالانه (mm)

WS: میانگین سرعت باد (m/s)

این متغیرها دامنه تغییرپذیری متفاوتی از تغییرات بسیار محدود در VPD تا نوسانات گسترده در بارش را نشان دادند؛ چنین پراکندگی‌ای

می‌توانست خطر بیش‌برازش را افزایش دهد. از این‌رو، این ویژگی‌ها از مدل‌سازی نهایی کنار گذاشته شده و تنها نقش کمکی در تفسیر تفاوت‌های مکانی بین مناطق داشته‌اند.

جمع‌بندی فرایند تهیه و آماده‌سازی داده‌ها

مراحل اصلی پژوهش، شامل گردآوری و پاک‌سازی داده‌ها، انتخاب مجموعه متغیرهای مؤثر، توسعه مدل‌های پیش‌بینی، ارزیابی عملکرد و تفسیر نتایج، در قالب یک فلوچارت تدوین شد تا روند تحقیق و قابلیت بازتولید نتایج روشن باشد. به‌منظور شفاف‌سازی روند اجرای پژوهش حاضر، مراحل اصلی از گردآوری داده‌ها تا تحلیل نهایی در قالب یک فلوچارت ارائه شده است. این نمودار به‌طور خلاصه فرایند تحقیق شامل آماده‌سازی داده‌ها، انتخاب ویژگی‌های مؤثر، توسعه مدل‌های پیش‌بینی، ارزیابی عملکرد مدل‌ها و تفسیر نتایج نهایی را نشان می‌دهد.

همان‌طور که در فلوچارت (شکل ۲) نشان داده شده است، تمامی مراحل تحقیق به‌صورت سیستماتیک و مرحله‌به‌مرحله دنبال شد و ساختار مقاله نیز بر همین اساس سازماندهی شده است.

پیش‌پردازش داده‌ها

داده‌های مورد‌استفاده در این پژوهش شامل متغیرهای اقلیمی، شاخص‌های سنجش‌ازدور و مقادیر عملکرد نهایی دانه گندم بودند که بر اساس سال زراعی و موقعیت مکانی سازمان‌دهی شدند. در مرحله نخست، داده‌ها از فایل اکسل وارد محیط برنامه‌نویسی شده و پس از بررسی اولیه، نام ستون‌ها یکسان‌سازی و ویژگی‌های موردنیاز برای مدل‌سازی انتخاب گردید. متغیر هدف مدل، عملکرد نهایی دانه (Yield) بوده و متغیرهای ورودی شامل شاخص NDVI و متغیرهای اقلیمی VPD، Tmax، Tmin، WS، SR و prcp بودند. همچنین برای افزایش پایداری مدل‌ها، نسبت عملکرد واقعی به عملکرد مرجع به‌عنوان متغیر هدف در فرایند آموزش مدل‌ها مورد‌استفاده قرار گرفت و عملکرد نهایی با استفاده از مقادیر مرجع بازیابی شد.

برای تحلیل حساسیت مدل‌ها و بررسی نقش عامل‌های اقلیمی در تبیین عملکرد گندم مناسب است. این شش متغیر، مجموعه اصلی ورودی‌های اقلیمی در مدل‌سازی نهایی را تشکیل دادند.

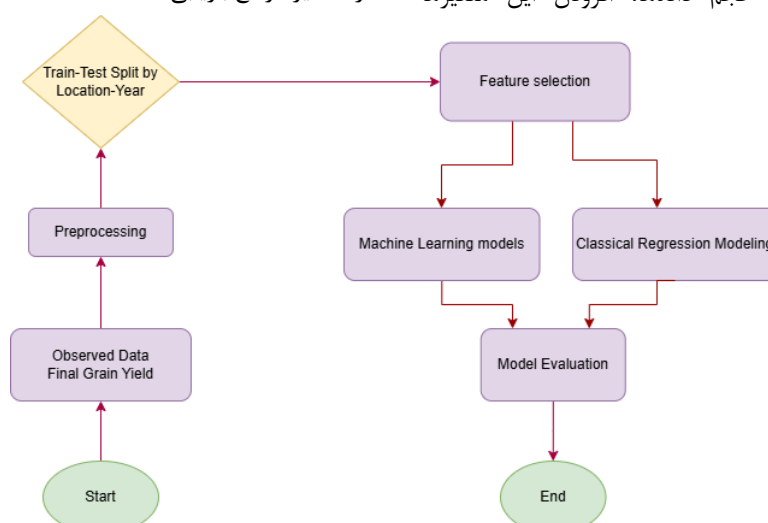
داده‌های سنجش‌ازدور و شاخص‌های پوشش گیاهی

در سال‌های اخیر، ترکیب داده‌های سنجش‌ازدور ماهواره‌ای و الگوریتم‌های یادگیری ماشین به یکی از رویکردهای اصلی در سامانه‌های کشاورزی دقیق تبدیل شده است. شاخص‌های پوشش گیاهی مشتق‌شده از تصاویر ماهواره‌ای مانند NDVI اطلاعات ارزشمندی درباره وضعیت سلامت و رشد گیاه فراهم می‌کنند و در بسیاری از مطالعات به‌عنوان متغیرهای کلیدی در مدل‌های پیش‌بینی عملکرد به کار گرفته شده‌اند [۳۶، ۳۷]. ترکیب این شاخص‌ها با متغیرهای اقلیمی و خاکی در چارچوب الگوریتم‌های یادگیری ماشین امکان استخراج الگوهای پیچیده رشد محصول را فراهم می‌کند.

از میان شاخص‌های مرتبط با وضعیت پوشش گیاهی، شاخص NDVI به‌عنوان متغیر کلیدی بیانگر سبزی‌نگی و فعالیت فتوسنتزی گندم مورد‌استفاده قرار گرفت. اگرچه داده‌های LAI نیز در مجموعه اولیه موجود بود، بررسی‌های توصیفی نشان داد مقدار آن در تمام رکوردها ثابت است؛ بنابراین در فرایند مدل‌سازی استفاده نشد. مقادیر NDVI که از قبل برای هر منطقه - سال محاسبه شده بود، بدون نیاز به پردازش تصاویر خام ماهواره‌ای به‌صورت مستقیم به مدل‌ها وارد شد.

داده‌های خاک و سایر ویژگی‌ها

ویژگی‌های متعددی مرتبط با خاک (شامل عمق خاک، ظرفیت مزرعه، نقطه پژمردگی، محتوای آب قابل‌دسترس، چگالی ظاهری، هدایت هیدرولیکی اشباع، درصد رس، درصد شن، کربن آلی خاک و شاخص‌های مرتبط با رشد ریشه) در نسخه اولیه داده‌ها وجود داشت. اما نتایج تحلیل مقدماتی نشان داد که بسیاری از این متغیرها یا تغییرپذیری مکانی اندک دارند یا همبستگی ضعیفی با عملکرد نشان می‌دهند. همچنین باتوجه به حجم داده‌ها، افزودن این متغیرها



شکل ۲: فلوچارت مراحل انجام تحقیق حاضر. مراحل شامل جمع‌آوری و آماده‌سازی داده‌ها، انتخاب ویژگی‌ها، توسعه مدل‌های پیش‌بینی، ارزیابی عملکرد مدل‌ها و تحلیل نتایج نهایی است.

Fig. 2: Flowchart of the present study procedure. The stages include data collection and preprocessing, feature selection, development of predictive models, evaluation of model performance, and analysis of the final results.

دقیق‌تر اهمیت واقعی هر ویژگی با استفاده از روش‌های تکمیلی نظیر Permutation Importance فراهم کرد.

مدل‌های کلاسیک رگرسیون

در گام نخست، مجموعه‌ای از توابع رگرسیون تک‌متغیره شامل خطی، نمایی، لگاریتمی، توانی و چندجمله‌ای درجه دوم، برای تحلیل رابطه بین هر ویژگی به صورت جداگانه با عملکرد نهایی محصول (Final Grain Yield) بررسی گردید. این مدل‌ها فقط یک متغیر مستقل را در تحلیل دخیل می‌کنند و برای بررسی تأثیر هم‌زمان چند ویژگی، باید از مدل‌های چندمتغیره مانند رگرسیون خطی چندگانه یا الگوریتم‌های یادگیری ماشین و یادگیری عمیق استفاده شود.

در این مرحله، برای ارزیابی توان پیش‌بینی ترکیبی از ویژگی‌های اقلیمی منتخب (طبق تحلیل Feature Importance) شامل تابش خورشیدی (SR)، بارش (prcp)، دمای بیشینه (Tmax)، دمای کمینه (Tmin)، کمبود فشار بخار (VPD) و سرعت باد (WS)، چهار روش رگرسیون چندمتغیره مورد استفاده قرار گرفت. این روش‌ها شامل رگرسیون خطی معمولی (Multiple Linear Regression)، رگرسیون چندجمله‌ای درجه دوم (Polynomial Regression - deg ۲) و دو روش منظم‌سازی شده Ridge Regression و Lasso Regression بودند. همه مدل‌ها روی داده‌های آزمون آموزش داده شدند و سپس عملکرد آن‌ها روی داده‌های تست ارزیابی شد. معیارهای ارزیابی شامل ریشه میانگین مربعات خطا (RMSE)، میانگین قدرمطلق خطا (MAE) و ضریب تعیین (R^2) در هر دو مجموعه آموزش و آزمون محاسبه شدند. این تحلیل به انتخاب مناسب‌ترین مدل و ارزیابی اهمیت ترکیبی متغیرها در پیش‌بینی عملکرد نهایی محصول کمک می‌کند.

الگوریتم‌های یادگیری ماشین

در مرحله بعد، مدل‌های پیچیده‌تر یادگیری ماشین از جمله (Random Forest, XGBoost, Gradient Boosting) به کار گرفته شدند. داده‌ها در قالب ویژگی (Features) و برچسب (Target) به مدل‌ها داده شدند و تنظیمات بهینه مدل‌ها از طریق جستجوی شبکه‌ای (Grid Search) با اعتبارسنجی متقابل ۵-تایی انجام شد. عملکرد مدل‌ها بر اساس مجموعه آزمون با استفاده از شاخص‌های RMSE و MAE و ضریب تعیین R^2 ارزیابی شد.

تحلیل اهمیت ویژگی‌ها

برای شناسایی میزان تأثیر هر یک از ویژگی‌های اقلیمی بر پیش‌بینی عملکرد نهایی محصول، روش Permutation Importance به کار گرفته شد. در این روش، مقادیر هر ویژگی به طور تصادفی جابه‌جا می‌شوند و کاهش دقت مدل اندازه‌گیری می‌شود. ویژگی‌هایی که بیشترین کاهش دقت را موجب شوند، بیشترین اهمیت پیش‌بینی را دارند. این روش به ویژه برای شناسایی متغیرهایی که همبستگی بالا با دیگر ویژگی‌ها دارند؛ اما اثر پیش‌بینی‌کننده مستقل دارند، مفید است. با استفاده از این روش، مجموعه‌ای از ویژگی‌های کلیدی برای آموزش مدل‌های

به منظور مدیریت مقادیر گمشده، از روش جایگزینی با میانه (Median Imputation) با استفاده از الگوریتم SimpleImputer استفاده شد. این روش با حفظ ساختار آماری داده‌ها، از حذف نمونه‌ها و کاهش حجم داده جلوگیری می‌کند.

باتوجه به تفاوت ماهیت مدل‌های مورد استفاده، پیش‌پردازش داده‌ها به صورت اختصاصی برای هر مدل انجام شد. برای مدل Ridge Regression، به منظور حذف اثر اختلاف مقیاس متغیرها و بهبود فرایند بهینه‌سازی، از روش استانداردسازی (StandardScaler) استفاده شد؛ به گونه‌ای که هر ویژگی به میانگین صفر و انحراف معیار یک تبدیل گردید. در مقابل، برای مدل‌های Random Forest و XGBoost استانداردسازی اعمال نشد، زیرا این مدل‌های مبتنی بر درخت تصمیم نسبت به مقیاس متغیرهای ورودی حساس نبوده و عملکرد آن‌ها تحت تأثیر اختلاف مقیاس ویژگی‌ها قرار نمی‌گیرد.

به منظور ارزیابی عملکرد مدل‌ها و جلوگیری از نشت اطلاعات، از اعتبارسنجی متقاطع گروهی پنج‌تایی (GroupKFold, 5-fold Cross-Validation) استفاده شد. در این روش، سال زراعی به عنوان متغیر گروه‌بندی در نظر گرفته شد؛ به گونه‌ای که تمامی نمونه‌های مربوط به هر سال تنها در یکی از فولدهای آموزش یا اعتبارسنجی قرار گرفتند. در هر تکرار، مدل با استفاده از چهار فولد آموزش داده شده و عملکرد آن بر روی فولد باقی‌مانده ارزیابی گردید. این رویکرد علاوه بر جلوگیری از انتقال اطلاعات بین سال‌های زراعی، ارزیابی واقع‌بینانه‌تری از قابلیت تعمیم مدل‌ها در پیش‌بینی داده‌های سال‌های جدید فراهم می‌کند.

در ادامه، برای درک بهتر نقش هر متغیر در تغییرات عملکرد محصول، تحلیل همبستگی بین ویژگی‌ها و عملکرد نهایی انجام شد.

تحلیل همبستگی بین ویژگی‌ها و عملکرد نهایی محصول

به منظور شناسایی روابط بین ویژگی‌های ورودی و متغیر هدف (عملکرد نهایی محصول) و همچنین بررسی احتمال هم‌خطی بین متغیرها، تحلیل همبستگی انجام شد. در این پژوهش، از ضریب همبستگی پیرسون (Pearson Correlation Coefficient) برای محاسبه میزان وابستگی خطی بین هر ویژگی اقلیمی و خاکی با عملکرد محصول استفاده گردید. همچنین، برای شناسایی روابط بین خود ویژگی‌ها، ماتریس همبستگی بین تمام متغیرهای ورودی محاسبه و به صورت تصویری با heatmap نمایش داده شد. این تحلیل دو هدف اصلی داشت:

- شناسایی ویژگی‌هایی که بیشترین ارتباط را با عملکرد محصول دارند تا بتوانند در مراحل انتخاب ویژگی و توسعه مدل به عنوان متغیرهای کلیدی لحاظ شوند.

- شناسایی ویژگی‌هایی که همبستگی بالایی با یکدیگر دارند تا از وقوع هم‌خطی (Multicollinearity) جلوگیری شود و پیچیدگی مدل کاهش یابد.

نتایج این تحلیل به عنوان پایه‌ای برای انتخاب ویژگی‌های مؤثر در مدل‌سازی بعدی مورد استفاده قرار گرفت و راهنمایی برای بررسی

الگوهای غیرخطی و تشخیص مقادیر دورافتاده را فراهم می‌کنند. همچنین با استفاده از ضرایب آماری مانند R^2 و شاخص‌های خطا شامل RMSE و MAE که بر روی نمودارها نمایش داده شده‌اند، امکان تفسیر کمی دقت مدل وجود دارد. مجموعه این ابزارها چارچوبی مناسب برای ارزیابی جامع عملکرد مدل‌های XGBoost، Random Forest و Ridge فراهم می‌کند.

روش آموزش و اعتبارسنجی مدل‌ها

برای ارزیابی عملکرد مدل‌های رگرسیون و یادگیری ماشین مورد استفاده در این پژوهش، از روش اعتبارسنجی متقاطع گروهی پنج‌تایی (GroupKFold, 5-fold Cross-Validation) استفاده شد. در این روش، سال زراعی (Season) به عنوان متغیر گروه‌بندی در نظر گرفته شد؛ به گونه‌ای که تمامی نمونه‌های مربوط به هر سال تنها در یکی از فولدهای آموزش یا اعتبارسنجی قرار گرفتند. در هر تکرار، مدل با استفاده از چهار فولد آموزش داده‌شده و عملکرد آن بر روی فولد باقی‌مانده ارزیابی گردید. این فرایند در پنج تکرار انجام شد، به طوری که تمامی نمونه‌ها دقیقاً یکبار در مجموعه اعتبارسنجی قرار گرفتند. استفاده از GroupKFold موجب جلوگیری از انتقال اطلاعات بین سال‌های زراعی و ارائه ارزیابی واقع‌بینانه‌تری از قابلیت تعمیم مدل در پیش‌بینی داده‌های سال‌های جدید شد.

ارزیابی عملکرد مدل و معیارهای آماری

به منظور تعیین مقادیر بهینه ابرپارامترهای مدل‌های یادگیری ماشین، از روش جست‌وجوی شبکه‌ای جامع (Grid Search) در چارچوب اعتبارسنجی متقاطع گروهی پنج‌تایی (GroupKFold, 5-fold Cross-Validation) استفاده شد. در این روش، سال زراعی (Season) به عنوان متغیر گروه‌بندی در نظر گرفته شد؛ به گونه‌ای که تمامی نمونه‌های مربوط به هر سال به طور کامل تنها در یکی از مجموعه‌های آموزش یا اعتبارسنجی قرار گیرند. این رویکرد از انتقال اطلاعات بین داده‌های آموزش و اعتبارسنجی جلوگیری کرده و امکان ارزیابی واقع‌بینانه‌تری از توانایی تعمیم مدل در پیش‌بینی سال‌های جدید را فراهم می‌کند.

برای هر یک از مدل‌های Ridge Regression، Random Forest و XGBoost، مجموعه‌ای از مقادیر ممکن برای ابرپارامترهای اصلی تعریف شد و تمامی ترکیب‌های آن‌ها با استفاده از روش Grid Search مورد ارزیابی قرار گرفت. معیار انتخاب بهترین ترکیب، بیشترین میانگین ضریب تعیین (R^2) در فولدهای اعتبارسنجی بود. به منظور افزایش قابلیت تعمیم مدل‌ها و کاهش احتمال بیش‌برازش، تنظیم ابرپارامترها بر اساس نتایج اعتبارسنجی متقاطع انجام شد. در مدل XGBoost ابرپارامترهای `learning_rate`، `max_depth`، `n_estimators`، `reg_lambda` و `subsample` در مدل Random Forest پارامترهای `n_estimators`، `min_samples_leaf` و `max_features` و در مدل Ridge Regression پارامتر α (alpha) مورد بررسی قرار گرفتند. مقادیر نهایی ابرپارامترهای انتخاب‌شده در جدول ۱ ارائه شده است.

پیش‌بینی انتخاب شد تا دقت پیش‌بینی به حداکثر برسد و پیچیدگی مدل کاهش یابد.

ارزیابی مدل‌های یادگیری ماشین با ترکیب ویژگی‌ها

در مرحله بعد، با هدف شناسایی بهترین ترکیب ویژگی‌ها برای پیش‌بینی عملکرد نهایی محصول، تمامی زیرمجموعه‌های ممکن از شش ویژگی اقلیمی انتخاب‌شده با استفاده از تحلیل Feature Importance و Permutation Importance بررسی شدند. برای هر ترکیب، تمامی مدل‌ها آموزش داده‌شده و عملکرد آن‌ها با معیارهای RMSE، MAE و R^2 ارزیابی شد. این رویکرد اجازه می‌دهد تا اثر ترکیبی متغیرها بر دقت پیش‌بینی مشخص شود و ویژگی‌هایی که بیشترین نقش را در بهبود عملکرد مدل دارند، شناسایی شوند.

اعتبارسنجی متقابل پنج‌تایی

جهت ارزیابی دقت، پایداری و توان تعمیم‌پذیری مدل‌های پیش‌بینی شده، از روش اعتبارسنجی متقابل پنج‌تایی (5-fold cross-validation) استفاده شد. در این روش، مجموعه داده‌های آموزش به پنج بخش (fold) مساوی تقسیم شد؛ به طوری که در هر تکرار، چهار بخش برای آموزش مدل و یک بخش باقی‌مانده برای آزمون موقت (validation) استفاده شد. این فرایند پنج بار تکرار شد تا هر بخش به عنوان داده آزمون موقت یکبار مورد استفاده قرار گیرد.

با اجرای این روش، مقادیر RMSE، MAE و R^2 برای هر fold محاسبه و سپس میانگین آن‌ها به عنوان شاخص عملکرد کلی مدل گزارش شد. این رویکرد برای تمام مدل‌های مورد استفاده اعمال شد تا عملکرد واقعی مدل‌ها در شرایط مختلف داده‌های آموزش به طور جامع مورد سنجش قرار گیرد.

تحلیل نمودارهای باقی‌مانده و دقت پیش‌بینی مدل‌ها

برای ارزیابی رفتار خطاهای مدل و بررسی کیفیت برازش، از نمودارهای توزیع باقی‌مانده‌ها استفاده شد. باقی‌مانده‌ها به صورت تفاضل بین مقادیر مشاهده‌شده و پیش‌بینی‌شده محاسبه شدند و توزیع آن‌ها می‌تواند اطلاعات ارزشمندی در مورد بایاس مدل، نرمال بودن خطاها و وجود الگوهای سیستماتیک ارائه دهد. نمودارهای هیستوگرام همراه با تابع چگالی احتمال (Kernel Density Estimation) به منظور بررسی شکل کلی توزیع خطا، پهنای توزیع، تجمع خطا در اطراف مقدار صفر و وجود مقادیر پرت (outliers) ترسیم شدند. تمرکز اصلی این تحلیل، ارزیابی این بود که آیا مدل‌ها تمایل به بیش‌برآوردی یا کم‌برآوردی دارند و اینکه آیا الگوی مشخصی در خطاها مشاهده می‌شود یا خیر؛ زیرا وجود الگو می‌تواند نشان‌دهنده نقص ساختاری در مدل‌سازی باشد.

برای سنجش دقت کلی پیش‌بینی مدل‌ها نیز از نمودارهای Observed vs Predicted استفاده شد. در این نمودارها، مقادیر پیش‌بینی‌شده در برابر مقادیر واقعی ترسیم شده و خط ۱:۱ (خط همانی) به عنوان مبنا برای ارزیابی کیفیت برازش مدل به کار رفت. هرچه نقاط بیشتر در نزدیکی این خط قرار گیرند، مدل از دقت بالاتری برخوردار است. این نمودارها امکان ارزیابی بصری میزان پراکندگی پیش‌بینی‌ها، شناسایی

در فرمول‌های فوق:

n : تعداد نمونه‌ها.

y_i : مقدار مشاهده‌شده

$\hat{y}_i$: مقدار پیش‌بینی‌شده توسط مدل

$\bar{y}$: میانگین مقادیر مشاهده‌شده می‌باشد.

نتایج و بحث

الگوی توزیع مکانی عملکرد گندم

بر اساس نقشه نشان‌داده‌شده در شکل ۳، عملکرد گندم زمستانه در آلمان دارای الگوی کاملاً ناهمگن است. بخش‌های شمالی کشور - به‌ویژه مناطق Schleswig و Rostock - به‌طور واضح در کلاس‌های عملکرد بالا قرار دارند. این نواحی در نقشه با رنگ‌های سبز پررنگ مشخص شده‌اند و گستره وسیعی از عملکردهای بالاتر از میانگین را نشان می‌دهند.

در مقابل، مناطق جنوب و جنوب غرب کشور مانند Freudenstadt و Augsburg با رنگ‌های زرد و سبز کم‌رنگ در نقشه مشخص شده‌اند و حاکی از عملکرد کمتر هستند. این الگو بیانگر وجود یک شیب مکانی شمال - جنوب است؛ جایی که افزایش دما، تبخیر - ترق و کاهش رطوبت نسبی می‌تواند نقش اصلی را در کاهش عملکرد ایفا کند.

تحلیل آماری و پراکندگی داده‌ها

در شکل ۴، در داده‌های اقلیمی و عملکردی ارائه‌شده برای نه منطقه مختلف، الگوهای مشخصی از تفاوت‌های مکانی در عملکرد نهایی محصول دیده می‌شود. نمودار جعبه‌ای نشان می‌دهد که مناطق Rostock و Schleswig معمولاً دارای مقادیر میانه بالاتر و دامنه پراکندگی نسبتاً محدودتر هستند؛ امری که بیانگر پتانسیل تولیدی بیشتر و ثبات بالاتر عملکرد در این مکان‌هاست. در مقابل، مناطقی نظیر Augsburg و برخی نواحی جنوبی‌تر با مقادیر میانه پایین‌تر مشخص می‌شوند که احتمالاً بازتاب حساسیت بیشتر آن‌ها به نوسانات اقلیمی سالانه یا محدودیت‌های خاکی است. وجود نقاط پرت در چندین منطقه نیز نشان‌دهنده وقوع سال‌هایی با شرایط ویژه از نظر دما، تبخیر و ترق، باد یا بارندگی است که موجب افزایش یا کاهش غیرمعمول عملکرد شده است.

این الگوها با داده‌های اقلیمی و خاکی ثبت‌شده، سازگار هستند و نشان می‌دهند که عملکرد نهایی حاصل برهم‌کنش پیچیده عوامل اقلیمی نظیر حداکثر و حداقل دما (Tmax, Tmin)، فشار بخار اشباع (VPD)، سرعت باد (WS)، تابش خورشیدی (SR) و بارندگی (prcp) و نیز ویژگی‌های خاکی از جمله ظرفیت نگهداشت آب (AWC)، بافت خاک، چگالی ظاهری (SBD) و کربن آلی (SOC) است. تفاوت بین مناطق با میانه‌های بالا و نوسان کم در برابر مناطقی با پراکندگی زیاد، نشان‌دهنده نقش تعیین‌کننده شرایط محیطی پایدارتر و ظرفیت بالاتر خاک در تعدیل تنش‌های اقلیمی است.

جدول ۱: ابرپارامترهای نهایی انتخاب‌شده برای مدل‌های Ridge Regression، Random Forest و XGBoost در پیش‌بینی عملکرد گندم زمستانه.

Table 1: Final hyperparameters selected for Ridge Regression, Random Forest, and XGBoost models for winter wheat yield prediction.

مدل‌ها Models	ابزارهای نهایی Final hyperparameters
Ridge Regression	$\alpha=10$
Random Forest	min_samples_leaf=2, n_estimators=1500 max_features=sqrt
XGBoost	learning_rate=0.05, max_depth=4, n_estimators=900 reg_lambda=5, colsample_bytree=0.9, subsample=0.9

بررسی نتایج جست‌وجوی شبکه‌ای نشان داد که عملکرد مدل XGBoost نسبت به تغییرات محدود در ابرپارامترهای اصلی از پایداری مناسبی برخوردار است. به‌طور مشخص، افزایش تعداد درخت‌ها از ۹۰۰ به ۱۱۰۰ بهبود معناداری در مقدار R^2 ایجاد نکرد و تنها موجب افزایش هزینه محاسباتی شد. همچنین افزایش عمق درخت‌ها بیش از ۴ باعث کاهش عملکرد مدل در اعتبارسنجی متقاطع گردید که نشان‌دهنده آغاز بیش‌برازش بود. از سوی دیگر، نرخ یادگیری ۰.۰۵ بهترین تعادل را میان دقت پیش‌بینی و قابلیت تعمیم مدل فراهم کرد؛ درحالی‌که مقادیر بزرگ‌تر موجب همگرایی سریع‌تر اما کاهش عملکرد مدل شدند. برای سنجش عملکرد مدل‌های توسعه‌یافته، از سه معیار آماری استاندارد و کاربرد استفاده شده است که در ادامه به تعریف و فرمول آن‌ها می‌پردازیم.

○ ریشه میانگین مربعات خطا (RMSE)

این معیار، میزان انحراف مقادیر پیش‌بینی‌شده از مقادیر واقعی را اندازه‌گیری می‌کند. فرمول آن به‌صورت زیر است:

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (\hat{y}_i - y_i)^2}$$

○ میانگین خطای مطلق (MAE)

این معیار، میانگین خطاهای مطلق بین مقادیر پیش‌بینی‌شده و واقعی را نشان می‌دهد و به‌دقت مدل اشاره دارد. فرمول آن به‌صورت زیر تعریف می‌شود:

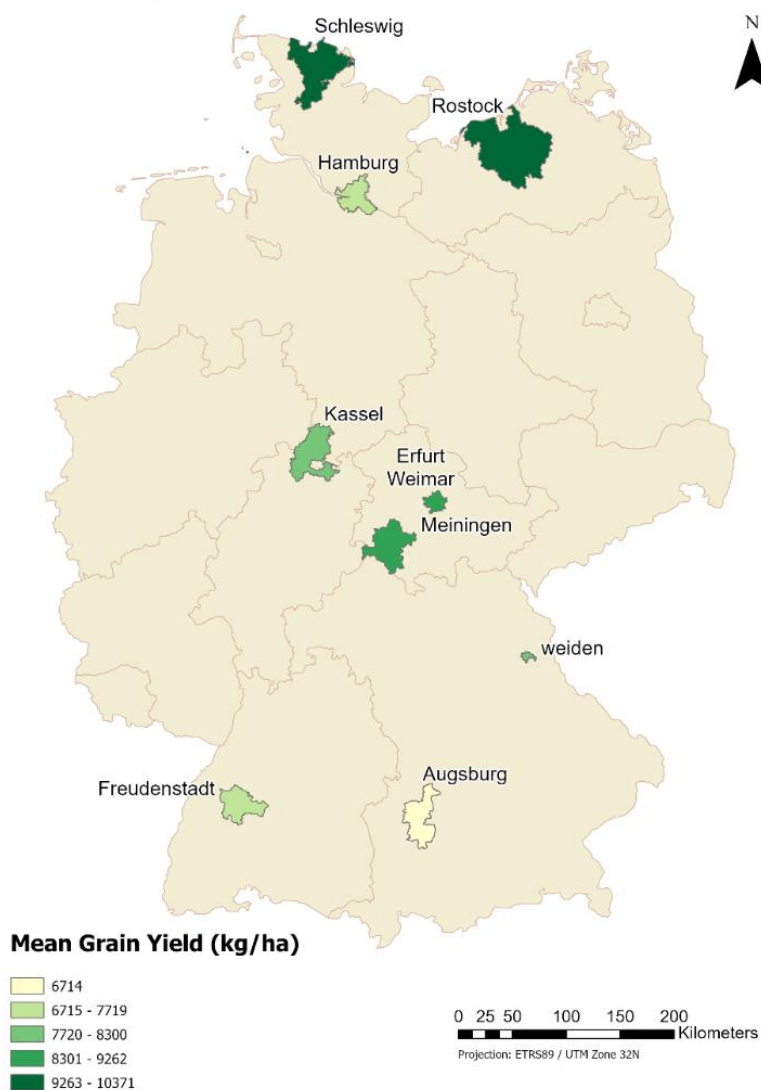
$$MAE = \frac{1}{n} \sum_{i=1}^n |\hat{y}_i - y_i|$$

○ ضریب تعیین (R^2)

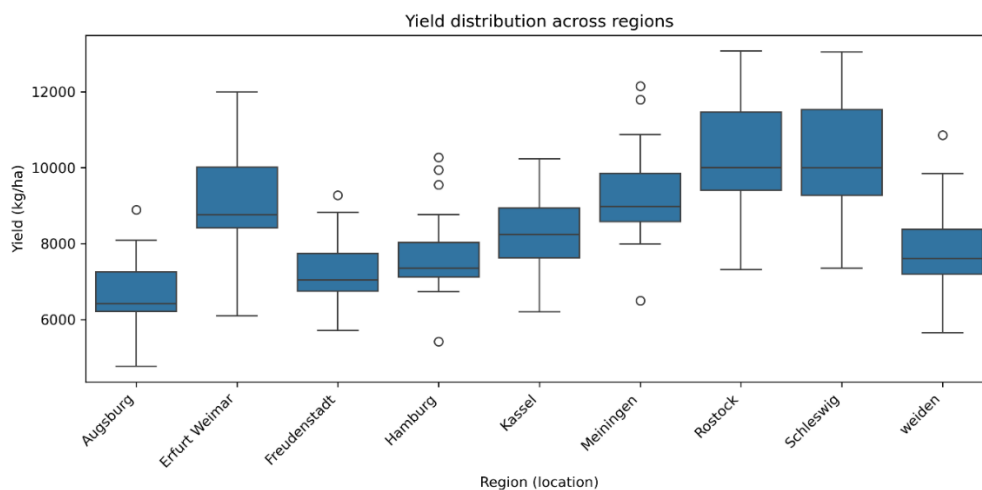
این معیار، میزان موفقیت مدل در توضیح واریانس داده‌های واقعی را می‌سنجد. هرچه مقدار آن به ۱ نزدیک‌تر باشد، نشان‌دهنده تطابق بهتر مدل با داده‌ها است. فرمول آن به‌صورت زیر است:

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2}$$

Spatial Pattern of Mean Wheat Yield



شکل ۳: الگوی توزیع مکانی میانگین عملکرد گندم زمستانه (بر حسب کیلوگرم در هکتار) در ۹ منطقه مورد مطالعه آلمان.
Fig. 3: Spatial distribution pattern of the average winter wheat yield (kg/ha) in the nine study regions of Germany.



شکل ۴: نمودار جعبه‌ای عملکرد گندم زمستانه در ۹ منطقه مورد مطالعه آلمان.
Fig. 4: Boxplot of winter wheat yield in the nine study regions of Germany.

این الگو با نقش فیزیولوژیک Tmin در کاهش تنش سرمایی و افزایش فعالیت فتوسنتزی و نقش باد در تهویه و کاهش رطوبت برگ در شرایط غیر تنش‌زا سازگار است. همبستگی منفی عملکرد با SR (تابش خورشیدی) نیز می‌تواند نشان‌دهنده شرایطی باشد که تابش بالا همراه با دما و تبخیر زیاد رخ داده و منجر به تنش گرمایی یا خشکی شده است. از سوی دیگر، NDVI همبستگی مثبت اما متوسطی با عملکرد دارد (۰٫۲۲) که نشان می‌دهد شاخص پوشش سبزیگی به‌تنهایی قادر به توضیح کامل تغییرات عملکرد نیست؛ اما همچنان نشانه‌ای از وضعیت رشد و زیست‌توده (Biomass) گیاهی فراهم می‌کند. همچنین روابط بین متغیرهای اقلیمی نیز الگوهای قابل‌انتظاری را نشان می‌دهند: Tmax و Tmin همبستگی مثبت بالا دارند (۰٫۶۴) و VPD با SR و WS نیز ارتباط مثبت نشان می‌دهد که همگی بازتاب سازوکارهای فیزیکی تبخیر و انرژی ورودی هستند. مجموعه این الگوها نشان می‌دهد که عملکرد نهایی محصول در این سامانه داده‌ای عمدتاً تابع شرایط رطوبتی و تعادل انرژی محیط است.

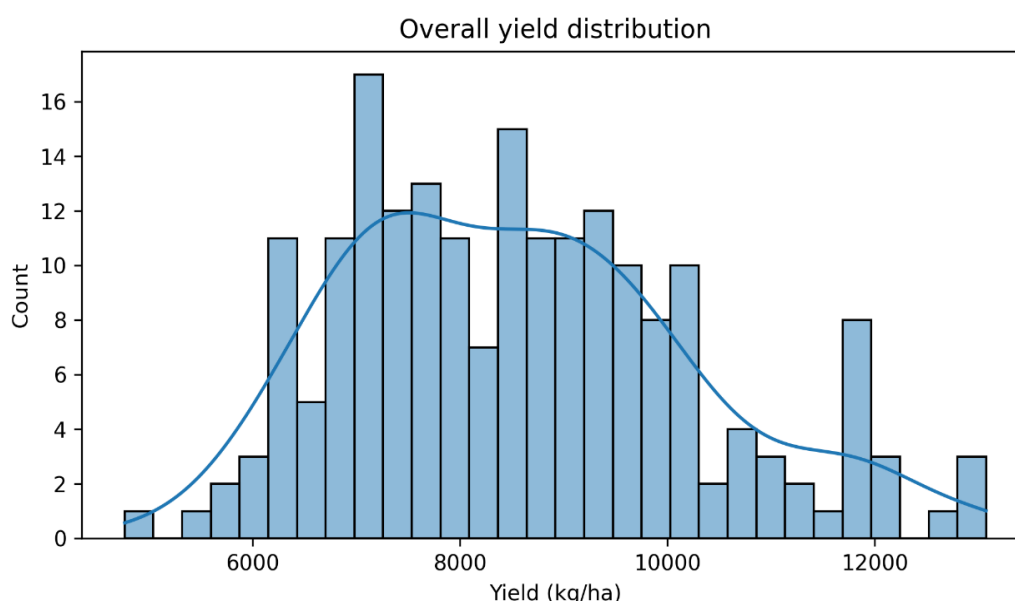
ارزیابی عملکرد مدل‌ها

○ مقایسه عملکرد مدل‌ها در پیش‌بینی (OOE (Out-of-Fold) تحلیل نمودارهای پراکندگی (شکل‌های ۷ تا ۹) نشان می‌دهد که مدل XGBoost عملکرد دقیق‌تر و پایدارتری نسبت به Ridge Regression و Random Forest دارد. در شکل مربوط به XGBoost، نقاط پیش‌بینی‌شده به‌صورت فشرده و نزدیک به خط ۴۵ درجه قرار گرفته‌اند؛ وضعیتی که نشان‌دهنده نبود سوگیری سیستماتیک و برازش مناسب در بازه‌های عملکرد پایین تا متوسط است که بیشترین حجم داده‌ها را تشکیل می‌دهند.

شکل ۵، هیستوگرام توزیع کلی عملکرد نشان می‌دهد که مقادیر Yield در مجموعه داده دارای الگوی تقریباً دو‌نمایی (bimodal) یا دست‌کم پهنای توزیع گسترده با یک قله اصلی در بازه ۷۰۰۰ تا ۹۰۰۰ کیلوگرم در هکتار است. بخش قابل‌توجهی از سال‌ها و مناطق در این محدوده قرار دارند که نشان‌دهنده شرایط عمومی اقلیمی و خاکی نسبتاً مناسب در بیشتر نقاط داده است. درعین حال، کشیدگی توزیع به سمت راست (right-skewed tail) وجود تعداد قابل‌توجهی مقادیر بالاتر از ۱۱۰۰۰ کیلوگرم نشان می‌دهد که برخی مناطق یا سال‌ها دارای شرایط استثنائاً مطلوب بوده‌اند؛ شرایطی که احتمالاً با ترکیبی از بارندگی مناسب، دماهای متوسط، ظرفیت نگهداشت آب بالاتر خاک همراه بوده است. از سوی دیگر، حضور تعداد محدود اما مشخص از سال‌ها با عملکرد کمتر از ۶۰۰۰ کیلوگرم بیانگر وقوع دوره‌های تنش‌زا مانند خشکی، دمای بالا، یا محدودیت رطوبتی خاک است؛ بنابراین این هیستوگرام علاوه بر توصیف شکل کلی توزیع عملکرد، بازتاب‌دهنده ناهمگنی قابل‌توجه در شرایط اقلیمی و خاکی است و اهمیت لحاظ کردن این عوامل را در مدل‌سازی یا تفسیرهای مبتنی بر سنجش از دور برجسته می‌کند.

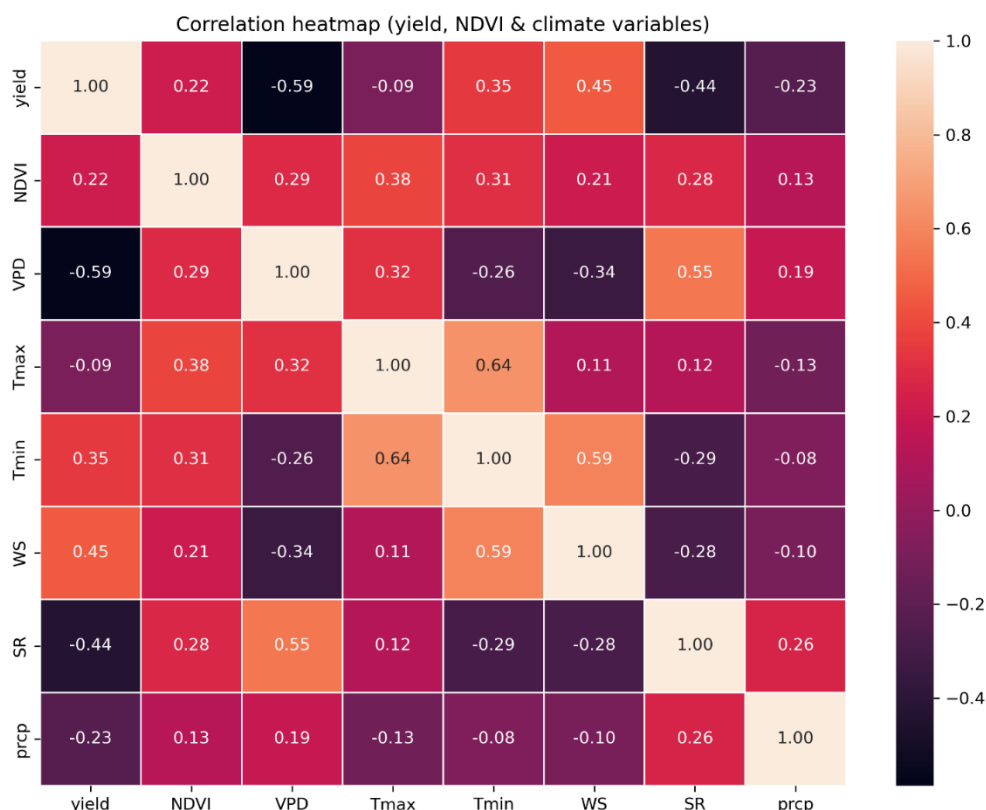
همبستگی ویژگی‌ها با عملکرد

نقشه همبستگی ارائه‌شده نشان می‌دهد که عملکرد نهایی محصول بیش از همه تحت‌تأثیر متغیرهای رطوبتی و تنش تبخیری قرار دارد. همبستگی منفی نسبتاً قوی بین عملکرد و VPD (حدود -۰٫۵۹) بیانگر این است که افزایش کسری فشار بخار - که نشانه تبخیر و ترقق بالقوه بالا و محدودیت رطوبتی گیاه است - به‌طور قابل‌توجهی عملکرد را کاهش می‌دهد. در مقابل، سرعت باد (WS) و حداقل دما (Tmin) همبستگی مثبت با عملکرد دارند، به‌طوری که WS با ضریب تقریباً ۰٫۳۵ و Tmin با حدود ۰٫۴۵ بیشترین ارتباط مثبت را نشان می‌دهند.



شکل ۵: هیستوگرام توزیع کلی عملکرد گندم زمستانه در تمام مناطق مورد مطالعه.

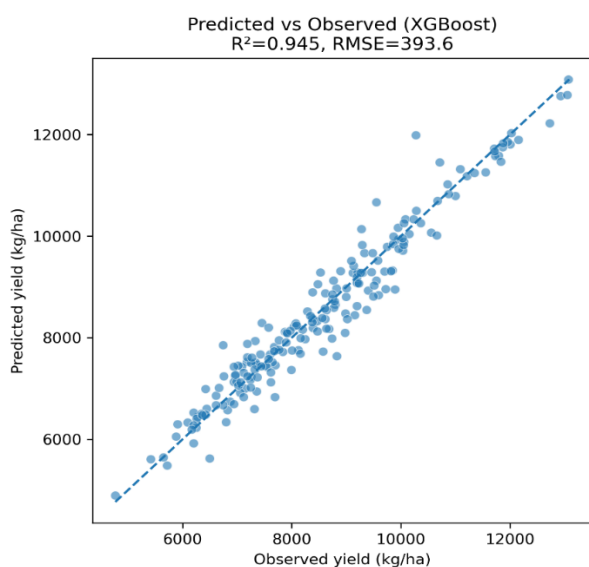
Fig. 5: Histogram of the overall distribution of winter wheat yield across all study regions.



شکل ۶: ماتریس همبستگی پیرسون بین عملکرد گندم، NDVI و متغیرهای کلیدی اقلیمی (VPD، Tmax، Tmin، WS، SR و بارش).

Fig. 6: Pearson correlation matrix between wheat yield, NDVI, and key climatic variables (VPD, Tmax, Tmin, WS, SR, and precipitation).

است که اندازه خطاهای این مدل نسبت به دو مدل دیگر به‌طور محسوسی کوچک‌تر است. علاوه بر این، مقایسه MAE در جدول ۲ نشان می‌دهد که XGBoost کمترین خطای مطلق میانگین را ثبت کرده و در نتیجه، عملکرد آن در کاهش خطا در کل دامنه مقادیر مشاهده‌شده بهتر است.



شکل ۷: مقایسه مقادیر مشاهده‌شده و پیش‌بینی‌شده عملکرد گندم با استفاده از مدل XGBoost

Fig. 7: Comparison of observed and predicted wheat yield values using the XGBoost model.

در مقابل، مدل Random Forest پراکندگی بیشتری را نشان می‌دهد و در مقادیر بالا دچار کم‌برآوردی می‌شود؛ رفتاری که معمولاً در مدل‌های مبتنی بر درخت (در صورت محدودیت عمق یا رزولوشن تقسیمات) مشاهده می‌شود. مدل Ridge Regression نیز بیشترین انحراف از خط ایدئال را دارد، به‌طوری که در عملکردهای بالا کم‌برآوردی و در عملکردهای پایین بیش‌برآوردی مشاهده می‌شود؛ نتیجه‌ای که با ماهیت خطی مدل و ناتوانی آن در بازنمایی روابط غیرخطی میان ورودی‌ها و عملکرد محصول سازگار است. مجموعه این الگوها بیانگر آن است که XGBoost دقیق‌ترین و پایدارترین پیش‌بینی OOF را ارائه می‌دهد و از هر دو مدل دیگر از نظر خطا، همگامی با خط ایدئال و ثبات در کل دامنه مقادیر برتر است.

○ مقایسه شاخص‌های R²، RMSE و MAE

نتایج سه شاخص اصلی ارزیابی OOF که در جدول ۲ ارائه شده‌اند، اختلاف عملکرد میان مدل‌ها را با وضوح بیشتری نشان می‌دهند. مطابق جدول ۲، مدل XGBoost بالاترین مقدار R² را به دست آورده است؛ بدین معنا که بیشترین توانایی را در توضیح واریانس عملکرد گندم دارد. مقدار R² مدل Random Forest کمی کمتر است و Ridge پایین‌ترین مقدار را ثبت کرده است که این امر بیانگر محدودیت مدل خطی Ridge در بازنمایی الگوهای پیچیده و غیرخطی میان متغیرهاست.

الگوی مشاهده‌شده در شاخص R² دقیقاً در شاخص‌های خطا نیز تکرار می‌شود. XGBoost کمترین مقدار RMSE را دارد؛ موضوعی که بیانگر آن

مدل گزینه مناسب‌تری نسبت به Ridge و Random Forest Regression برای استفاده به‌عنوان مدل نهایی به نظر می‌رسد.

جدول ۲: ارزیابی عملکرد مدل‌های یادگیری ماشین در پیش‌بینی عملکرد محصول بر اساس شاخص‌های R^2 ، RMSE و MAE.

Table 1: Evaluation of the performance of machine learning models in crop yield prediction based on R^2 , RMSE, and MAE indices.

مدل Model	ضریب تعیین R^2	ریشه میانگین مربعات خطا RMSE	میانگین قدرمطلق خطا MAE
XGBoost	0.94	393.64	287.18
Random Forest	0.92	468.44	392.42
Ridge Regression	0.88	559.07	453.87

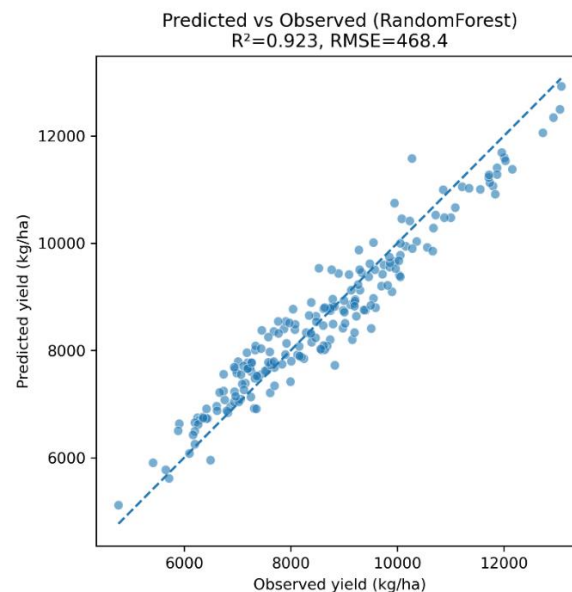
* توانایی مدل‌ها در بازتولید الگوهای مکانی

مقایسه نقشه واقعی عملکرد (شکل ۱۰-الف) با نقشه‌های خروجی مدل‌ها (شکل‌های ۱۰-ب تا ۱۰-د) نشان می‌دهد که XGBoost تنها مدلی است که قادر است الگوی مکانی عملکرد گندم را با دقت بالا و وضوح قابل قبول بازتولید کند. نقشه XGBoost در شکل ۱۰-ب نه تنها الگوی کلی شمال - جنوب را حفظ می‌کند، بلکه تغییرات ریزمقیاس درون منطقه‌ای را نیز به‌درستی تشخیص می‌دهد و مناطق با عملکرد بالا را بدون صاف‌سازی بیش از حد نشان می‌دهد. در مقابل، Random Forest در شکل ۱۰-ج اگرچه ساختار کلی را بازتاب می‌دهد، اما کنتراست مکانی آن کمتر است و مدل در مناطقی با تنوع زیاد، الگو را صاف‌تر و با جزئیات کمتر بازنمایی می‌کند. Ridge Regression در شکل ۱۰-د بیشترین میزان ساده‌سازی مکانی را اعمال کرده و جزئیات مهم از جمله شدت اختلاف میان شمال و جنوب را کاهش داده است. این مقایسه مکانی به‌وضوح نشان می‌دهد که XGBoost نسبت به دو مدل دیگر قادر است نه تنها مقدار عملکرد، بلکه ساختار فضایی آن را نیز با دقت بسیار خوبی بازسازی کند.

تحلیل اهمیت متغیرها

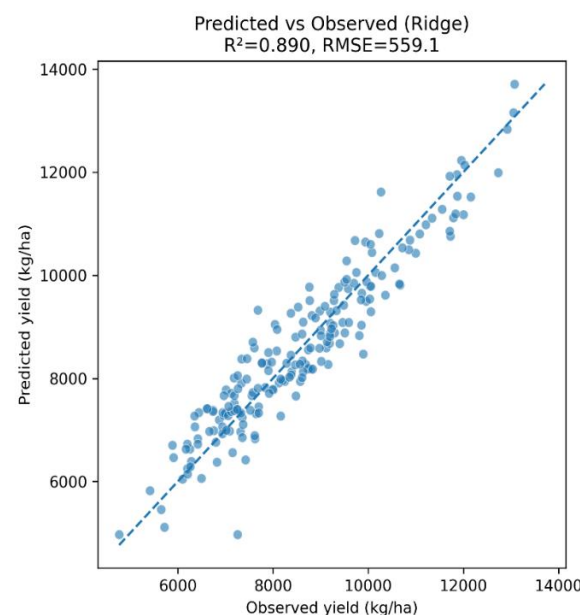
○ تحلیل اهمیت متغیرها در مدل XGBoost

بر اساس شکل ۱۱، مدل XGBoost نشان می‌دهد که VPD مهم‌ترین عامل در پیش‌بینی عملکرد گندم است و نسبت به سایر ویژگی‌ها بیشترین کاهش R^2 را در صورت جابه‌جایی دارد. این یافته بیانگر نقش بسیار پررنگ تنش رطوبتی و وضعیت تبخیر و تعرق در افت عملکرد است. پس از VPD، ویژگی‌های سرعت باد (WS) و NDVI در مراتب بعدی اهمیت قرار دارند. NDVI به‌عنوان شاخصی از وضعیت پوشش گیاهی و فتوسنتز، همچنان از عوامل کلیدی محسوب می‌شود، اما شدت اثر آن در مقایسه با VPD کمتر است. سایر متغیرها (T_{max} , T_{min} , $prcp$, SR) نقش فرعی دارند. این الگو بیانگر توانایی بالای XGBoost در آشکارسازی اثرات غیرخطی و ترکیبی متغیرهای اقلیمی است.



شکل ۸: مقایسه مقادیر مشاهده‌شده و پیش‌بینی‌شده عملکرد گندم با استفاده از مدل Random Forest

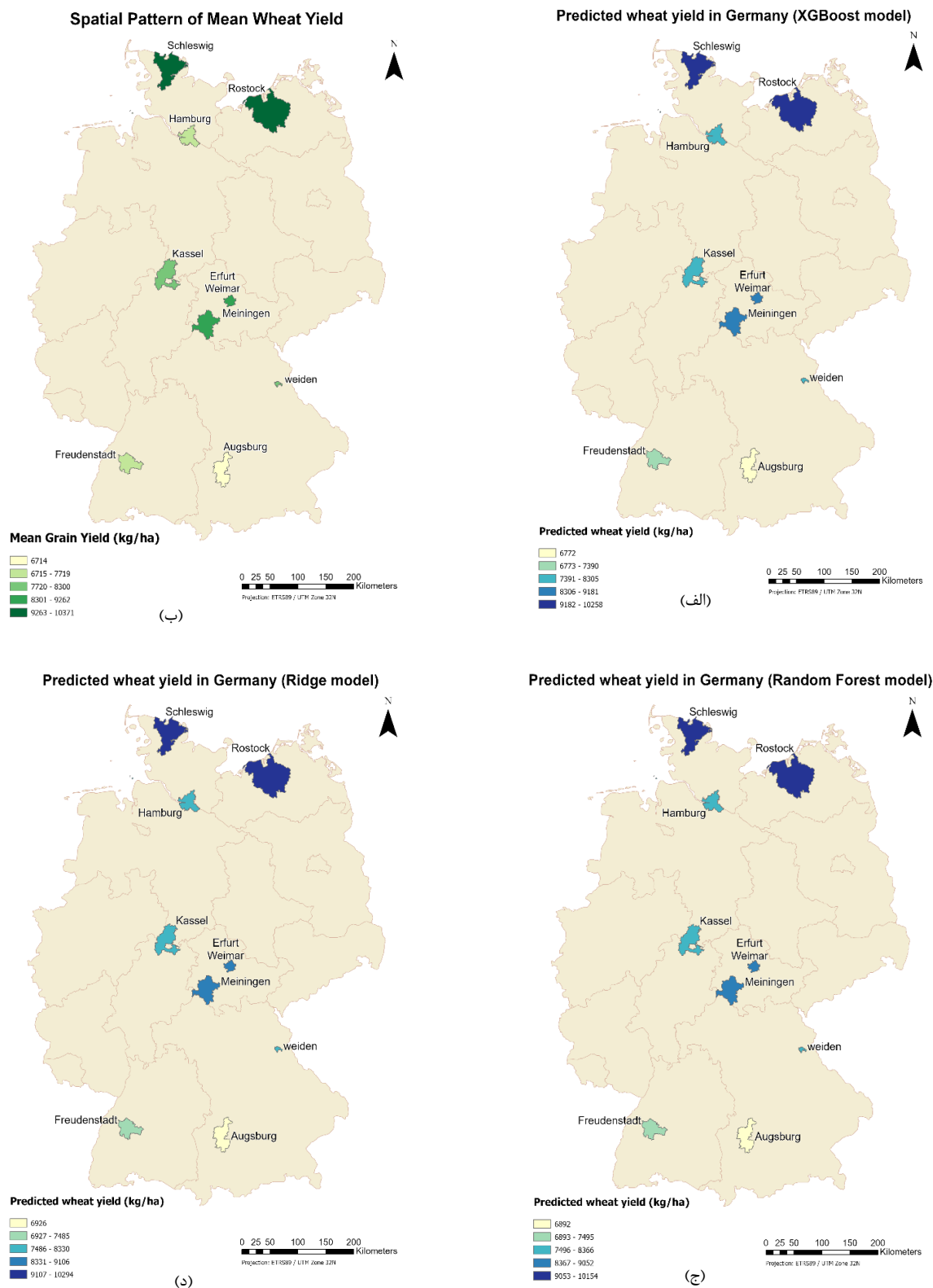
Fig. 8: Comparison of observed and predicted wheat yield values using the Random Forest model.



شکل ۹: مقایسه مقادیر مشاهده‌شده و پیش‌بینی‌شده عملکرد گندم با استفاده از مدل Ridge Regression

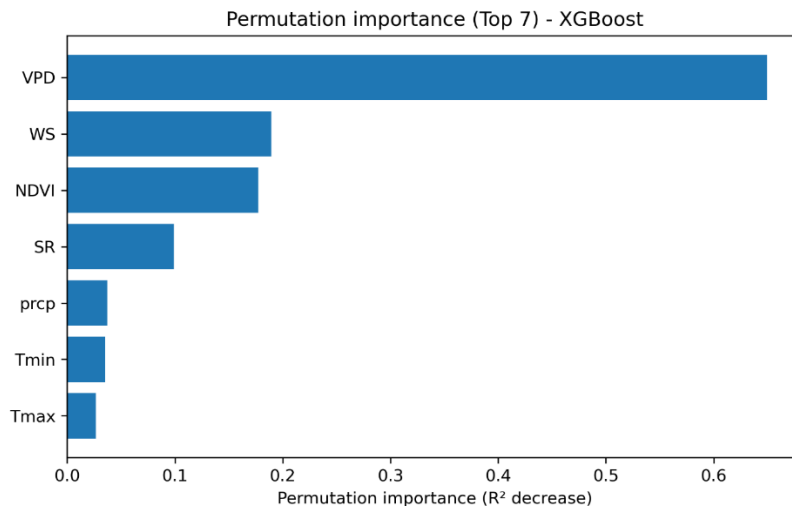
Fig. 9: Comparison of observed and predicted wheat yield values using the Ridge Regression model.

مجموع این سه شاخص R^2 بالا و خطاهای RMSE و MAE پایین - به‌صورت هم‌سو و سازگار تأیید می‌کنند که XGBoost در مسئله پیش‌بینی عملکرد که ماهیتی آشکارا غیرخطی دارد، بهترین توازن میان دقت، پایداری و توانایی تعمیم را ارائه داده است. به همین دلیل، این



شکل ۱۰: الگوی مکانی عملکرد واقعی گندم در مقایسه با عملکرد پیش‌بینی‌شده توسط سه مدل یادگیری ماشین (XGBoost، Ridge و Random Forest) در نقاط مختلف آلمان.

Fig. 10: Spatial pattern of actual wheat yield compared with yield predicted by three machine-learning models (XGBoost, Ridge Regression, and Random Forest) across different locations in Germany.



شکل ۱۱: اهمیت نسبی متغیرها در مدل XGBoost بر اساس روش Permutation Importance.
Fig. ۱۱: Relative importance of input variables in the XGBoost model based on the permutation importance method.

بوده و امکان ارزیابی تفصیلی حساسیت عملکرد نسبت به محرک‌های زیستی و اقلیمی را فراهم می‌کند.

در مورد NDVI (شکل ۱۲-الف)، مقادیر بالای SHAP در نواحی شمالی مانند Rostock و Hamburg نشان می‌دهد که وضعیت پوشش گیاهی در این مناطق نقش تعیین‌کننده‌تری در تبیین تغییرات عملکرد داشته است. این الگو احتمالاً بازتاب‌دهنده شرایط مدیریتی و فیزیولوژیکی مطلوب‌تر، تعادل بهتر آب - گیاه، و همبستگی قوی‌تر بین شاخص‌های سبزینه و مقادیر عملکرد در این مناطق است. در مقابل، در برخی مناطق جنوبی و غربی اهمیت NDVI کمتر مشاهده می‌شود که نشان می‌دهد در این نواحی عوامل اقلیمی دیگری - نظیر رطوبت هوا، تبخیر - تعرق، یا محدودیت‌های حرارتی - نقش پررنگ‌تری در تعیین عملکرد داشته‌اند. در خصوص VPD (شکل ۱۲-ب)، بیشترین مقادیر SHAP در مناطق Schleswig, Rostock و Augsburg مشاهده می‌شود. این الگو بیانگر آن است که در این نواحی طی فصل رشد، تنش تبخیر - تعرق و فشار بخار اشباع‌نشده هوا عامل غالب محدودکننده عملکرد بوده است. حضور اهمیت بالای VPD در هر دو بخش شمالی و جنوبی کشور نشان می‌دهد که حساسیت عملکرد نسبت به تنش رطوبتی رفتاری یکنواخت و خطی ندارد و احتمالاً تحت تأثیر هم‌زمان الگوی بارش، خصوصیات خاک و سازگاری ارقام گندم قرار می‌گیرد.

نتایج ارائه‌شده در جدول ۴ نشان می‌دهد که شدت تأثیر NDVI و VPD بر عملکرد گندم در میان مناطق مختلف کاملاً ناهمگون است. در میان متغیرهای مورد بررسی، VPD به‌طور کلی نقش پررنگ‌تری نسبت به NDVI داشته و در اغلب مناطق مقدار میانگین مطلق SHAP آن بالاتر است. بیشترین مقدار اثرگذاری VPD در منطقه Schleswig مشاهده شد (۱۱۷۴٫۹۳) که حاکی از حساسیت بالای عملکرد گندم به تنش رطوبتی در این منطقه است. به همین ترتیب، منطقه Rostock نیز مقدار بالایی از اثر VPD را نشان می‌دهد (۱۰۰۰٫۸۰) که این موضوع با شرایط اقلیمی شمال آلمان و نوسانات رطوبت جو سازگار است.

جدول ۳: اهمیت نسبی متغیرهای ورودی در مدل XGBoost بر اساس روش Permutation Importance

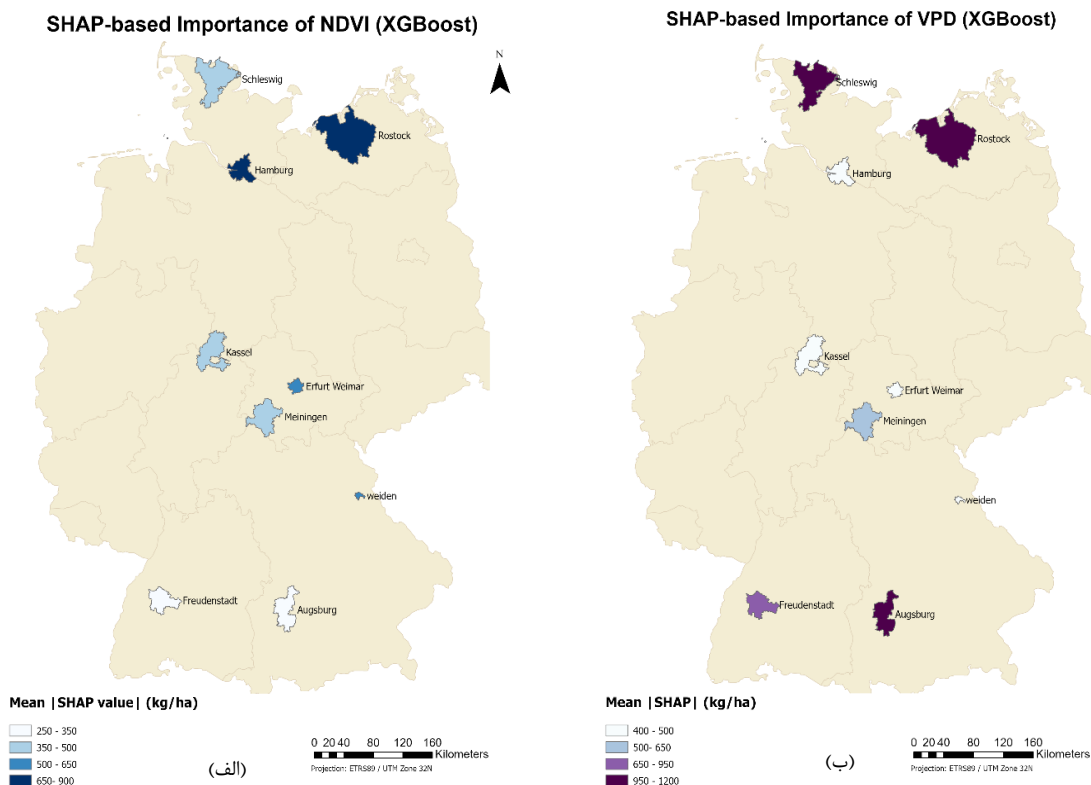
Table 2: Relative importance of input variables in the XGBoost model based on the Permutation Importance method.

متغیر	میانگین اهمیت	انحراف معیار
Variable	Mean Importance	Standard Deviation
VPD	0.65	0.049
WS	0.19	0.016
NDVI	0.18	0.013
SR	0.10	0.012
prcp	0.04	0.006
Tmin	0.04	0.004
Tmax	0.03	0.002

بر اساس نتایج ارائه‌شده در جدول ۳، متغیر VPD با مقدار میانگین اهمیت حدود ۰٫۶۵ بیشترین نقش را در پیش‌بینی عملکرد گندم در مدل XGBoost داشته است و فاصله قابل‌توجهی با سایر متغیرها دارد. پس از آن سرعت باد (WS) و NDVI به ترتیب با مقادیر اهمیت حدود ۰٫۱۹ و ۰٫۱۸ در رتبه‌های بعدی قرار می‌گیرند. این نتایج نشان می‌دهد که متغیرهای مرتبط با تنش تبخیر - تعرق و شرایط جوی نقش تعیین‌کننده‌تری در مدل داشته‌اند، درحالی‌که سایر متغیرهای اقلیمی مانند تابش خورشیدی (SR)، بارش (prcp)، حداقل و حداکثر دما (Tmin و Tmax)، سهم کمتری در توضیح تغییرات عملکرد داشته‌اند. مقدار انحراف معیار اهمیت متغیرها نیز نسبتاً پایین است که نشان‌دهنده پایداری نتایج اهمیت ویژگی‌ها در تکرارهای مختلف محاسبه است. در مجموع، این نتایج بیانگر آن است که مدل XGBoost قادر است تأثیر نسبی عوامل اقلیمی و زیستی را به‌خوبی تفکیک کرده و نقش غالب VPD به‌عنوان شاخصی از تنش رطوبتی را در تغییرات عملکرد گندم آشکار سازد.

○ تحلیل تکمیلی متغیرهای کلیدی NDVI و VPD

نتایج حاصل از تحلیل SHAP در شکل ۱۲ الگوی تغییرپذیری مکانی اهمیت دو متغیر کلیدی NDVI و VPD را در پیش‌بینی عملکرد گندم با استفاده از مدل XGBoost نشان می‌دهد. مقادیر SHAP در هر ناحیه بیانگر سهم واقعی و جهت اثرگذاری هر ویژگی در تولید خروجی مدل



شکل ۱۲: نقشه توزیع مکانی مقادیر SHAP مربوط به NDVI و VPD در مدل XGBoost.
Fig. 12: Spatial distribution maps of SHAP values for NDVI and VPD derived from the XGBoost model.

عملکرد گندم را نسبت به شرایط زیستی (NDVI) و اقلیمی (VPD) در مناطق مختلف تفکیک کرده و تأثیرپذیری محصول از تنش تبخیر - تعرق در شمال نسبت به جنوب کشور شدیدتر است.

در مجموع، تحلیل فضایی SHAP نشان می‌دهد که هرچند NDVI و VPD هر دو نقش کلیدی در عملکرد دارند، اما شدت و ماهیت اثرگذاری آن‌ها به‌طور محسوسی در نواحی مختلف متفاوت است. این ناهمگنی مکانی اهمیت ویژگی‌ها بر پیچیدگی روابط زیستی - اقلیمی حاکم بر عملکرد دلالت دارد و استفاده از مدل‌های غیرخطی و داده‌محور مانند XGBoost را برای استخراج الگوهای مکانی - فرایندی ضروری می‌سازد. چنین نتایجی بر ارزش افزوده رویکردهای مبتنی بر SHAP در تفسیرپذیری مدل‌های یادگیری ماشین در سیستم‌های کشاورزی تأکید می‌کند.

تحلیل باقیمانده‌ها و الگوهای خطا

○ تحلیل توزیعی باقیمانده‌ها

بررسی هیستوگرام‌های باقیمانده (شکل ۱۳) نشان می‌دهد که مدل XGBoost کمترین سوگیری سیستماتیک را ایجاد کرده و توزیع خطاهای آن تقریباً متقارن و متمرکز در حوالی صفر است. در این مدل، نه تنها میانگین خطا بسیار کوچک است، بلکه پراکندگی دوطرفه خطاها نیز نشان می‌دهد عملکرد بیش‌برآورد یا کم برآوردی قابل‌توجهی در بازه‌های مختلف وجود ندارد. در مقابل، توزیع خطا در Random Forest پهن‌تر است و تمایل به کم برآوردی در نمونه‌های با عملکرد بالا سبب

جدول ۴: مقادیر میانگین مطلق SHAP برای NDVI و VPD به تفکیک منطقه در مدل XGBoost.

Table 3: Mean absolute SHAP values for NDVI and VPD by region in the XGBoost model.

مکان Location	میانگین قدرمطلق مقدار NDVI برای SHAP SHAP_mean_abs_NDVI	میانگین قدرمطلق مقدار VPD برای SHAP SHAP_mean_abs_VPD
Augsburg	266.59	964.88
Erfurt Weimar	506.46	482.91
Freudenstadt	303.90	651.83
Hamburg	870.62	422.49
Kassel	401.64	494.07
Meiningen	494.34	517.55
Rostock	651.17	1000.80
Schleswig	491.82	1174.93
Weiden	538.99	462.13

در مقابل، NDVI اثر بیشتری در برخی مناطق دارد. بیشترین اهمیت NDVI در Hamburg دیده می‌شود (۸۷۰٫۶۲) که نشان‌دهنده نقش برجسته وضعیت پوشش گیاهی و سلامت محصول در تبیین عملکرد در این منطقه است. همچنین مناطق Rostock و Weiden نیز مقادیر نسبتاً بالای SHAP برای NDVI دارند (به ترتیب ۶۵۱٫۱۷ و ۵۳۸٫۹۹). این اختلافات مکانی بیانگر آن است که مدل XGBoost به‌دقت حساسیت

شکل‌دهی توزیع فضایی عملکرد است. نقشه پیش‌بینی نشان می‌دهد که مناطق شمالی و شمال شرق آلمان (به‌ویژه Rostock و Schleswig) بالاترین مقادیر عملکرد را دارند، درحالی‌که برخی از مناطق جنوب غربی و بخش‌هایی از شمال (مانند Hamburg و Freudenstadt) در طبقات عملکرد پایین تا متوسط قرار گرفته‌اند. این الگو به‌طور مستقیم با تفاوت‌های منطقه‌ای در اقلیم، خاک و پوشش گیاهی مرتبط است.

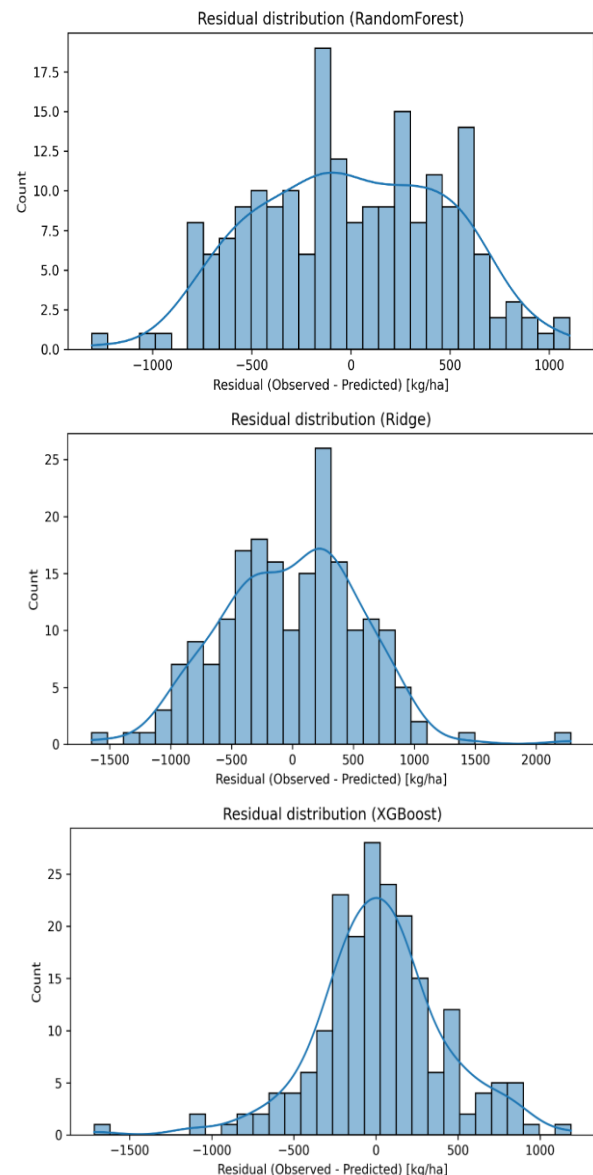
در مناطق پربازده مانند Rostock و Schleswig، ترکیب اقلیم دریایی ملایم (دمای معتدل، نوسان حرارتی کم، تنش تبخیری متوسط)، بارش کافی و تابش مناسب شرایطی را ایجاد کرده که تنش حرارتی و رطوبتی به حداقل رسیده و پایداری رشد در فصل زراعی تضمین می‌شود. این مناطق دارای پروفیل‌های خاکی عمیق‌تر، ظرفیت نگهداشت آب بالاتر و محدودیت کمتر رشد ریشه هستند و در نتیجه ریشه گیاه امکان دسترسی پایدار به رطوبت را دارد. همچنین مقادیر NDVI بالاتر (~۰.۲۲-۰.۲۳) نشان می‌دهد که کارایی فتوسنتزی در این مناطق مطلوب‌تر است. هم‌زمانی این شرایط باعث شده است که مناطق شمالی در طبقه عملکرد متوسط روبه‌بالا تا بالا قرار گیرند.

در مقابل، مناطق کم‌بازده یا متوسط روبه‌پایین نظیر Freudenstadt و Hamburg، هرچند از نظر دما و بارش در محدوده مطلوب قرار دارند، اما محدودیت‌های خاص اقلیمی یا خاکی ظرفیت تولید را کاهش داده است. در Freudenstadt، بارش بسیار زیاد همراه با دمای پایین منجر به کاهش تابش مؤثر و افزایش رطوبت خاک می‌شود که احتمال تهویه ضعیف خاک، محدودیت تنفس ریشه و بروز بیماری‌های قارچی را افزایش می‌دهد. در Hamburg، درحالی‌که اقلیم نسبتاً مناسب است، بافت شنی خاک، ماده آلی بسیار کم، AWC پایین و هدایت هیدرولیکی بالا موجب خروج سریع آب از ناحیه ریشه و کاهش کارایی جذب آب می‌شود. در این مناطق NDVI پایین‌تر (~۰.۱۵-۰.۱۸) نیز نشان‌دهنده رشد محدود پوشش گیاهی و کاهش تولید زیست‌توده است که با عملکرد پایین‌تر سازگار است.

در مناطق با عملکرد متوسط مانند Erfurt-Weimar، Meiningen، Kassel و Weiden، معمولاً یک تعادل نسبی بین شرایط مساعد و محدودکننده دیده می‌شود. این مناطق اغلب دارای اقلیم معتدل، بارش متوسط و تنش تبخیری پایین تا متوسط هستند، اما ویژگی‌های خاکی آنها (مانند AWC متوسط، ماده آلی متوسط یا محدودیت‌های ریشه‌ای) یا تابش کمتر نسبت به مناطق پربازده باعث شده عملکرد آنها در محدوده میانی قرار گیرد. مقادیر NDVI در این مناطق نیز در حد متوسط (~۰.۲۰-۰.۲۱) است که نشان‌دهنده توسعه کانوپی نه‌چندان محدود، اما به‌مراتب کمتر از مناطق پربازده شمالی است.

در مجموع، الگوی مکانی پیش‌بینی‌شده توسط مدل XGBoost نشان می‌دهد که بالاترین عملکردها در مناطق شمالی با اقلیم دریایی، خاک‌های عمیق و توسعه کانوپی بهتر رخ می‌دهد، درحالی‌که مناطق با بارش بسیار زیاد، تابش کم، یا خاک‌های سبک و کم‌عمق معمولاً در طبقات عملکرد پایین‌تر قرار می‌گیرند. این نتایج تأکید می‌کند که

شده دنباله سمت منفی کمی برجسته‌تر باشد. هیستوگرام Ridge نشان‌دهنده بیشترین عدم تقارن است و الگوی مشخصی از کم برآوردی برای عملکردهای بالا را نشان می‌دهد که ناشی از ناتوانی مدل خطی در مدل‌سازی روابط غیرخطی بین متغیرهاست. این تفاوت‌ها بیانگر آن است که XGBoost نه‌تنها از نظر برازش عددی بهتر عمل کرده، بلکه از منظر رفتار آماری خطاها نیز ساختاری پایدار و قابل‌اتکا دارد.



شکل ۱۳: توزیع فراوانی و چگالی باقیمانده‌ها در سه مدل رگرسیونی.
Fig. 13: Frequency distribution and density of residuals for the three regression models.

نقشه‌های پیش‌بینی عملکرد

○ الگوی مکانی عملکرد پیش‌بینی‌شده توسط XGBoost
نتایج پیش‌بینی مدل XGBoost الگوی مکانی مشخص و قابل تفسیری از عملکرد گندم در سراسر آلمان ارائه می‌کند که بیانگر نقش هم‌زمان عوامل اقلیمی، ویژگی‌های خاک و شاخص‌های پوشش گیاهی در

به‌طور کلی، این مقایسه نشان می‌دهد که مدل XGBoost قادر است الگوی فضایی عملکرد گندم را با دقت مناسبی بازسازی کند و مناطق با عملکرد بالا، متوسط و پایین را به‌خوبی تفکیک نماید. این نتایج بیانگر قابلیت مدل برای استفاده در تحلیل‌های مکانی عملکرد محصولات و برنامه‌ریزی مدیریتی در سیستم‌های کشاورزی است.

○ تحلیل مکانی خطاها در مدل XGBoost

شکل ۱۵ الگوی فضایی خطای مدل XGBoost را در مقیاس منطقه‌ای نشان می‌دهند. نقشه باقیمانده‌ها (Residuals) بیان می‌کند که بیشتر مناطق دارای خطاهای کوچک و فاقد سوگیری هستند و اختلاف بین عملکرد واقعی و پیش‌بینی‌شده در حد ± 100 کیلوگرم در هکتار باقی می‌ماند. تنها چند ناحیه، از جمله Hamburg و Schleswig، باقیمانده‌های بزرگ‌تری را نشان می‌دهند که می‌تواند ناشی از ناهمگنی شدید اقلیمی - خاکی یا تغییرات درون مزرعه‌ای باشد؛ عواملی که مدل و داده‌های ورودی سنجش‌ازدور همواره قادر به بازنمایی کامل آن‌ها نیستند.

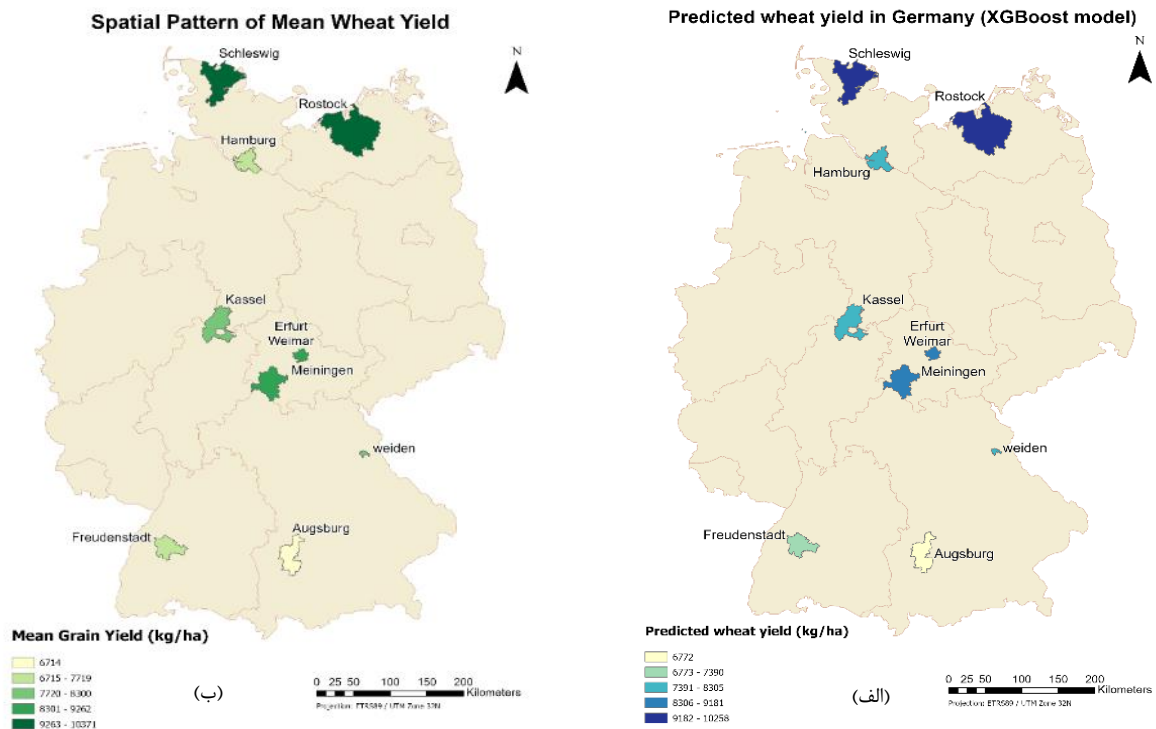
نقشه RMSE در شکل ۱۵ نشان می‌دهد که بیشترین میزان خطا در چند منطقه پراکنده شامل Hamburg، Meiningen و Freudenstadt مشاهده می‌شود. این مناطق معمولاً با شرایط اقلیمی متغیرتر، تنش رطوبتی متناوب یا ویژگی‌های خاکی ناهمگن همراه هستند این الگو با یافته‌های مطالعات پیشین همخوانی دارد و بیان می‌کند که مدل در محیط‌های باثبات عملکرد بسیار مطلوبی دارد.

عملکرد گندم در آلمان تابع برهم‌کنش پیچیده بین عوامل اقلیمی (دما، بارش، VPD، تابش)، خصوصیات خاک (عمق، AWC، SOC، RGF) و وضعیت رشد پوشش گیاهی (NDVI) است، و وضعیت مکانی عملکرد بیش از هر چیز توسط همین ترکیب منطقه‌ای شرایط محیطی شکل می‌گیرد.

○ مقایسه مکانی پیش‌بینی‌ها با عملکرد واقعی

مقایسه نقشه عملکرد واقعی گندم با نتایج پیش‌بینی‌شده توسط مدل XGBoost نشان می‌دهد که الگوی مکانی عملکرد در اغلب مناطق به‌خوبی توسط مدل بازتولید شده است. هر دو نقشه بیانگر عملکرد بالاتر در مناطق شمالی مانند Rostock و Schleswig و عملکرد پایین‌تر در برخی مناطق مانند Freudenstadt و Hamburg هستند. این هم‌خوانی نشان می‌دهد که مدل توانسته ارتباط میان عوامل اقلیمی، ویژگی‌های خاک و شاخص پوشش گیاهی (NDVI) با عملکرد گندم را به‌طور مؤثری شناسایی کند.

در مناطق با عملکرد متوسط مانند Erfurt-Weimar، Meiningen، Kassel و Weiden نیز تطابق قابل‌قبولی بین نتایج پیش‌بینی و داده‌های واقعی مشاهده می‌شود که نشان‌دهنده توانایی مدل در تشخیص شرایط محیطی میانی و تأثیر آن‌ها بر عملکرد است. با این حال، در برخی مناطق مانند Freudenstadt و Hamburg اختلافات جزئی بین عملکرد واقعی و پیش‌بینی‌شده دیده می‌شود که می‌تواند ناشی از ناهمگنی‌های محلی خاک یا شرایط اقلیمی خاص باشد که به‌طور کامل در داده‌های ورودی مدل منعکس نشده‌اند.



شکل ۱۴: الگوی فضایی عملکرد واقعی گندم و عملکرد پیش‌بینی‌شده توسط مدل XGBoost در آلمان.
Fig. 14: Spatial pattern of actual wheat yield and yield predicted by the XGBoost model in Germany.

خاک و الگوهای رشد نامنظم همراهاند. XGBoost در این شرایط نسبت به Ridge Regression و Random Forest پایداری بیشتری نشان داد، هرچند مقادیر RMSE همچنان بالاتر از میانگین باقی ماند. رفتار Random Forest به صورت کم برآوردی در مناطق پربازده نمایان شد و Ridge به دلیل ماهیت خطی، ناهمگنی فضایی را بیش از حد ساده‌سازی کرد؛ بنابراین، در چشم‌اندازهای اقلیمی متنوع، مدل‌های غیرخطی به‌ویژه XGBoost مزیت قابل توجهی در بازنمایی دقیق الگوهای مکانی دارند.

نتایج این بخش نشان می‌دهد که ساختار مکانی عملکرد گندم ارتباط نزدیکی با شرایط اقلیمی - زیستی دارد و مدل XGBoost بیشترین توانایی را در بازتاب این ساختار از خود نشان داده است. نقش برجسته NDVI و VPD در شکل‌گیری الگوهای فضایی، همراه با رفتار مکانی خطاها، بیانگر آن است که در مناطق گسترده و اقلیم‌های ناهمگن، استفاده از مدل‌های غیرخطی برای دستیابی به پیش‌بینی‌های دقیق ضروری است. XGBoost علاوه بر دقت عددی بالا، الگوهای فضایی عملکرد را نیز به صورت واقع‌گرایانه‌تری بازسازی می‌کند؛ ویژگی‌ای که آن را به گزینه‌ای بهینه برای مدل‌سازی مکانی عملکرد محصولات تبدیل می‌کند.

تفسیر و مقایسه یافته‌ها با مطالعات پیشین

هدف این پژوهش ارزیابی و مقایسه عملکرد مدل‌های مختلف یادگیری ماشین در پیش‌بینی عملکرد گندم با استفاده از داده‌های اقلیمی و شاخص سنجش‌ازدور NDVI در مناطق مورد مطالعه در آلمان بود. برای این منظور، مدل‌های XGBoost، Random Forest و Ridge Regression با استفاده از شاخص‌های آماری RMSE، MAE و ضریب تعیین (R^2) مورد ارزیابی قرار گرفتند تا توانایی آن‌ها در پیش‌بینی دقیق عملکرد و بازنمایی الگوهای مکانی تغییرات تولید بررسی شود.

نتایج نشان داد که مدل XGBoost در مقایسه با Random Forest و Ridge Regression عملکرد برتری در پیش‌بینی عملکرد گندم دارد. این برتری را می‌توان به توانایی بالای XGBoost در مدل‌سازی روابط غیرخطی و برهم‌کنش‌های پیچیده میان متغیرهای اقلیمی و عملکرد محصول نسبت داد. ساختار تقویتی این الگوریتم باعث می‌شود که خطاهای باقی‌مانده در هر مرحله به صورت تدریجی اصلاح شوند و مدل بتواند الگوهای پنهان در داده‌ها را با دقت بیشتری استخراج کند. این ویژگی در سامانه‌های کشاورزی که معمولاً با ناهمگنی مکانی و زمانی بالایی همراه هستند اهمیت ویژه‌ای دارد.

اگرچه Random Forest نیز قادر به ثبت بخشی از پیچیدگی‌های داده است، اما در این مطالعه عملکرد آن نسبت به XGBoost ضعیف‌تر بود. ماهیت تجمیعی این مدل و استقلال نسبی درخت‌ها باعث می‌شود که توانایی آن در بازنمایی برخی روابط ظریف و تغییرات مکانی پیوسته محدودتر باشد. در نتیجه، در برخی مناطق با تغییرپذیری اقلیمی بیشتر، دقت پیش‌بینی Random Forest کاهش یافته است. در مقابل، Ridge

شکل ۱۵ توزیع مکانی ضریب تعیین (R^2) را نمایش می‌دهد. مقادیر بالاتر R^2 در مناطق دارای تنوع کم اقلیمی، مانند Schleswig و Kassel، مشاهده می‌شود که نشان می‌دهد مدل XGBoost در این مناطق توانایی بیشتری برای استخراج الگوهای وابستگی میان متغیرها دارد. در مقابل، مقادیر پایین‌تر R^2 عمدتاً در مناطقی با نوسانات شدیدتر دما، رطوبت یا ویژگی‌های خاکی دیده می‌شود؛ شرایطی که روابط زیستی - اقلیمی را پیچیده‌تر کرده و می‌تواند از دقت مدل بکاهد.

در مجموع، تحلیل فضایی خطاها نشان می‌دهد که مدل XGBoost از نظر پایداری مکانی و نبود سوگیری سیستماتیک عملکرد بسیار خوبی دارد. بیشترین خطاها در مناطقی رخ می‌دهد که از نظر اقلیمی یا خاکی ناهمگن‌تر هستند، درحالی‌که مناطق پایدارتر به طور دقیق‌تری مدل‌سازی می‌شوند. این نتایج تأیید می‌کند که XGBoost نه تنها در شاخص‌های عددی، بلکه در رفتار فضایی خطا نیز نسبت به روش‌هایی مانند Random Forest و Ridge Regression عملکرد برتری دارد.

جمع‌بندی یکپارچه روندهای مکانی و رفتاری مدل‌ها

○ یکپارچه‌سازی الگوی مکانی عملکرد

نقشه‌های عملکرد واقعی نشان دادند که بیشترین مقادیر عملکرد در نواحی شمالی و شمال شرقی آلمان رخ می‌دهد، درحالی‌که بخش‌هایی از مرکز و جنوب غرب کشور عملکرد پایین‌تری ثبت می‌کنند. هر سه مدل مورد بررسی این الگوی کلی را بازتولید کردند، اما XGBoost دقیق‌ترین بازنمایی را از ساختار فضایی عملکرد ارائه داد. این موضوع تأیید می‌کند که الگوی مکانی عملکرد گندم به طور عمده حاصل تعامل میان وضعیت رطوبت، دما و شاخص‌های رشد گیاهی است. مدل‌های غیرخطی، به‌ویژه XGBoost، توانستند تغییرات ریزمقیاس مرتبط با مدیریت و اقلیم خرد را نیز با دقت بیشتری منعکس کنند.

○ نقش متغیرهای کلیدی در شکل‌دادن ساختار مکانی

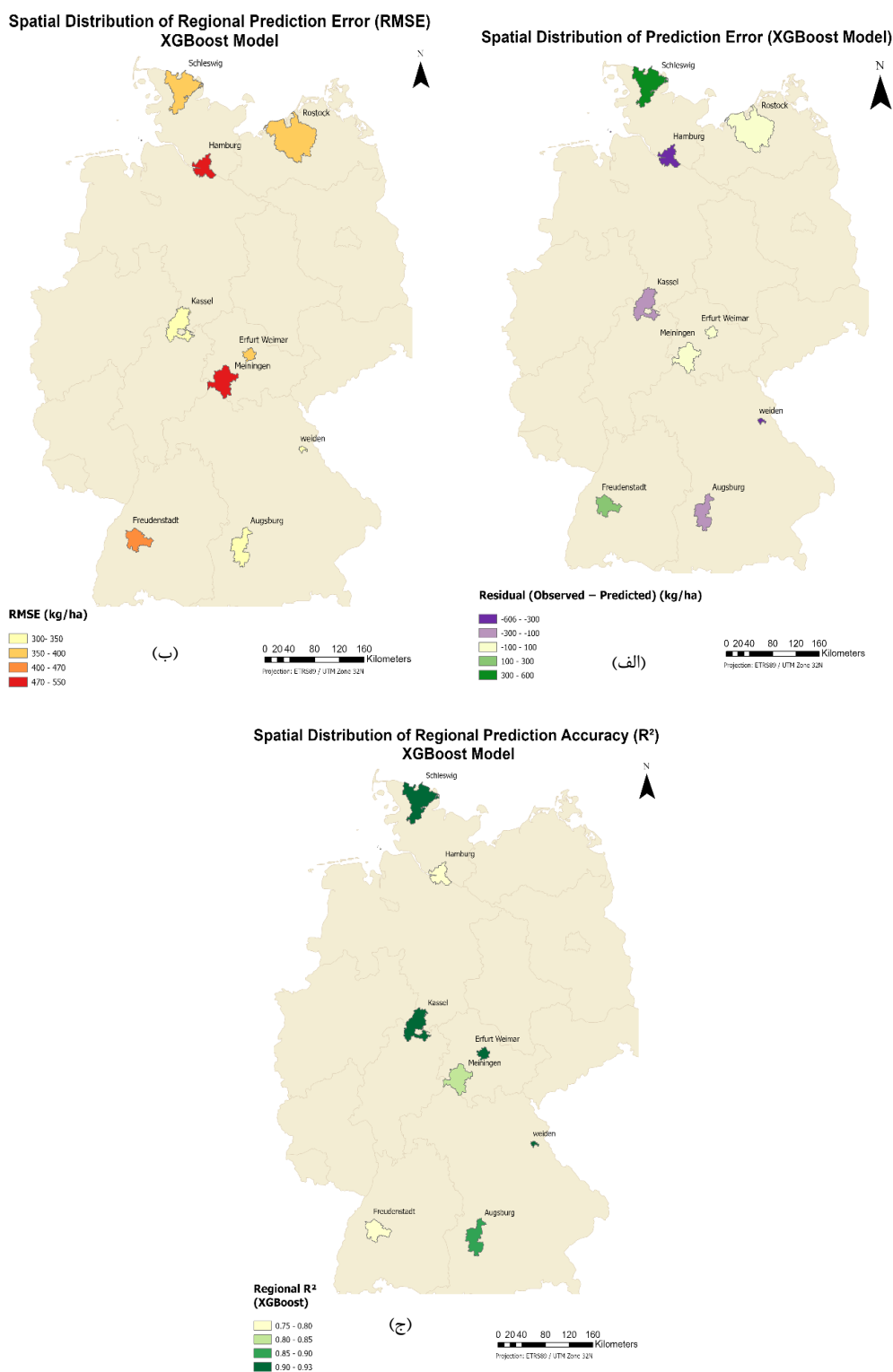
تحلیل اهمیت متغیرها نشان داد که NDVI و VPD نقش تعیین‌کننده‌ای در تبیین الگوی مکانی عملکرد دارند. NDVI به عنوان نمایه‌ای مستقیم از وضعیت زیستی و پویایی رشد، الگوی فضایی عملکرد را با دقت بازتاب می‌دهد. از سوی دیگر، VPD به عنوان شاخص فشار تبخیری، محدودیت‌های رطوبتی مؤثر بر تولید را مشخص می‌کند. ترکیب این دو متغیر سبب می‌شود مناطقی با NDVI بالا و VPD پایین - که عمدتاً در نواحی شمالی واقع‌اند - به بالاترین عملکرد دست یابند، درحالی‌که مناطق دارای NDVI پایین و VPD بالاتر عملکرد محدودتری نشان می‌دهند. این الگو در پیش‌بینی‌های XGBoost به صورت شفاف‌تر از سایر مدل‌ها مشاهده شد.

○ رفتار مدل‌ها در شرایط اقلیمی و خاکی متنوع

تحلیل خطاهای مکانی نشان داد که چالش اصلی مدل‌ها در مناطقی رخ می‌دهد که ناهمگنی اقلیمی و خاکی بیشتری دارند؛ از جمله Hamburg، Meiningen، Erfurt-Weimar و بخش‌هایی از جنوب غرب. این مناطق معمولاً با نوسانات زیاد دما و رطوبت، تغییرپذیری بالای خصوصیات

ناهمگنی واقعی سیستم‌های کشاورزی را نادیده بگیرد، به‌ویژه در مناطقی که تعاملات پیچیده میان عوامل اقلیمی و محیطی بر تولید محصول تأثیر می‌گذارند.

Regression به‌عنوان یک مدل خطی منظم‌سازی‌شده کمترین کارایی را نشان داد. این موضوع بیانگر آن است که ساده‌سازی روابط میان متغیرهای ورودی و عملکرد محصول می‌تواند بخش قابل‌توجهی از



شکل ۱۵: توزیع مکانی شاخص‌های خطای مدل XGBoost شامل: (الف) باقی‌مانده‌ها، (ب) RMSE و (ج) ضریب تعیین (R^2).
Fig. 15: Spatial distribution of XGBoost model error metrics: residuals, RMSE, and coefficient of determination (R^2).

و جهاجهاریا و همکاران (۲۰۲۳) نشان دادند که ترکیب داده‌های اقلیمی، خاک و سنجش‌ازدور می‌تواند دقت پیش‌بینی مدل‌های یادگیری ماشین را افزایش دهد. علاوه بر این، باسین و همکاران (۲۰۲۳) نشان دادند که ادغام یادگیری ماشین با داده‌های سنجش‌ازدور و سامانه‌های اینترنت اشیا می‌تواند نقش مهمی در توسعه سامانه‌های هوشمند مدیریت کشاورزی داشته باشد. درعین‌حال، برخی مطالعات نیز به چالش‌های این مدل‌ها اشاره کرده‌اند. به‌عنوان مثال، پائودل و همکاران (۲۰۲۳) بر اهمیت تفسیرپذیری مدل‌ها در کاربردهای عملی کشاورزی تأکید کرده‌اند و سابو و همکاران (۲۰۲۳) نشان داده‌اند که مدل‌های پیچیده‌تر در شرایط داده‌های محدود ممکن است با خطر بیش‌برازش مواجه شوند.

از نظر کاربردی، نتایج این پژوهش می‌تواند نقش مهمی در توسعه کشاورزی هوشمند و مدیریت پایدار تولید محصولات کشاورزی داشته باشد. استفاده از مدل‌های پیش‌بینی عملکرد مبتنی بر یادگیری ماشین امکان شناسایی دقیق‌تر الگوهای میان متغیرهای اقلیمی و تولید محصول را فراهم می‌کند و می‌تواند در توسعه سامانه‌های تصمیم‌یار کشاورزی مورداستفاده قرار گیرد. چنین سامانه‌هایی می‌توانند اطلاعات ارزشمندی برای برنامه‌ریزی زمان کاشت، انتخاب ارقام مناسب و مدیریت بهینه نهاده‌هایی مانند آب و کود فراهم کنند. علاوه بر این، پیش‌بینی عملکرد در مقیاس‌های منطقه‌ای یا ملی می‌تواند در برنامه‌ریزی‌های کلان کشاورزی، مدیریت بازار محصولات و ارزیابی وضعیت امنیت غذایی نیز مورداستفاده قرار گیرد.

نتیجه‌گیری

این پژوهش با هدف ارزیابی توانایی مدل‌های یادگیری ماشین در پیش‌بینی عملکرد گندم زمستانه در آلمان، با بهره‌گیری از داده‌های اقلیمی، شاخص سنجش‌ازدور (NDVI) و عملکرد نهایی محصول انجام شد. مقایسه سه رویکرد مدل‌سازی نشان داد که مدل XGBoost با ضریب تعیین ۰٫۹۴ و خطای RMSE برابر ۳۹۳٫۶ کیلوگرم بر هکتار، عملکرد دقیق‌تری نسبت به Random Forest و Ridge Regression دارد و توانسته الگوهای مکانی عملکرد را با خطای کمتر و قدرت تبیین بالاتر بازسازی کند. تحلیل اهمیت متغیرها نیز بیانگر نقش غالب تنش تبخیری (VPD) و به دنبال آن سرعت باد و شاخص NDVI در تعیین عملکرد بود؛ درحالی‌که متغیرهای دما و بارش سهم کمتری داشتند. تحلیل مکانی نشان داد که مدل XGBoost تفاوت‌های عملکردی میان مناطق شمالی (پربازده) و جنوبی (کم‌بازده) آلمان را با دقت مناسب بازتاب می‌دهد، هرچند در مناطق با ناهمگنی اقلیمی و خاکی بالا مانند Hamburg و Freudenstadt، خطای پیش‌بینی افزایش یافت. بااین‌حال، الگوی کلی عملکرد و ارتباط معنادار آن با شاخص‌های NDVI، رطوبت و تنش تبخیری به‌خوبی توسط مدل استخراج شد. به‌طورکلی، یافته‌ها نشان می‌دهد که مدل‌های تقویت گرادیان، به‌ویژه XGBoost، به دلیل توانایی در مدل‌سازی روابط غیرخطی و پیچیده میان عوامل محیطی،

بررسی شاخص‌های ارزیابی نیز این الگو را تأیید کرد. مدل XGBoost با دستیابی به مقادیر بالاتر R^2 و مقادیر کمتر RMSE و MAE نسبت به سایر مدل‌ها، توانایی بیشتری در توضیح تغییرات عملکرد محصول و کاهش خطای پیش‌بینی نشان داد. این نتایج نشان می‌دهد که مدل‌هایی که قادر به ثبت روابط غیرخطی و تعاملات پیچیده میان متغیرها هستند، می‌توانند پیش‌بینی دقیق‌تری از عملکرد محصولات کشاورزی ارائه دهند.

تحلیل نقشه‌های مکانی پیش‌بینی عملکرد و شاخص‌های خطا نیز دیدگاه مهمی درباره رفتار فضایی مدل‌ها فراهم کرد. نقشه‌های پیش‌بینی نشان دادند که مدل XGBoost توانسته است الگوی فضایی تغییرات عملکرد گندم را با انسجام و واقع‌گرایی بیشتری بازسازی کند و ناهمگنی‌های مکانی را نسبت به مدل‌های دیگر با جزئیات بیشتری نشان دهد. در مقابل، الگوی حاصل از Ridge هموارتر و ساده‌تر بود و بخشی از تغییرات ریزمقیاس عملکرد را منعکس نمی‌کرد که این موضوع با ماهیت خطی این مدل قابل توضیح است.

الگوی مکانی خطاها نیز نشان داد که میزان خطای پیش‌بینی در برخی مناطق بیشتر از سایر نواحی است. این تفاوت‌ها را می‌توان تا حد زیادی با ناهمگنی شرایط اقلیمی و محیطی توضیح داد. مناطقی که با نوسانات بیشتر دما و رطوبت، تغییرپذیری بارش یا تفاوت‌های قابل‌توجه در ویژگی‌های خاک مواجه هستند معمولاً دارای تغییرات عملکرد پیچیده‌تری هستند و در نتیجه پیش‌بینی آن‌ها برای مدل دشوارتر است. در مقابل، در مناطقی با شرایط اقلیمی پایدارتر و یکنواخت‌تر، الگوی عملکرد منظم‌تر بوده و مدل‌ها قادرند با دقت بیشتری تغییرات تولید را پیش‌بینی کنند. بررسی نقشه باقیمانده‌ها نیز نشان داد که در بیشتر مناطق خطاهای پیش‌بینی کوچک بوده و سوگیری سیستماتیک قابل‌توجهی در مدل مشاهده نمی‌شود.

یافته‌های این پژوهش با بخش قابل‌توجهی از مطالعات پیشین در زمینه کاربرد یادگیری ماشین در پیش‌بینی عملکرد محصولات کشاورزی همخوانی دارد. بسیاری از مطالعات نشان داده‌اند که الگوریتم‌های یادگیری ماشین قادرند روابط پیچیده میان متغیرهای اقلیمی، محیطی و عملکرد محصول را به‌خوبی مدل‌سازی کنند و دقت پیش‌بینی را نسبت به روش‌های آماری سنتی افزایش دهند. به‌عنوان مثال، شارما و همکاران (۲۰۲۵) نشان دادند که مدل‌های یادگیری ماشین، به‌ویژه روش‌های تجمعی مانند جنگل تصادفی و تقویت گرادیان، توانایی بالایی در پیش‌بینی عملکرد محصولات دارند. همچنین مرور نظام‌مند شاون و همکاران (۲۰۲۴) نشان داد که الگوریتم‌هایی مانند رگرسیون خطی، جنگل تصادفی و تقویت گرادیان از رایج‌ترین روش‌های مورداستفاده در این حوزه هستند و شاخص‌هایی مانند R^2 ، RMSE و MAE مهم‌ترین معیارهای ارزیابی عملکرد مدل‌ها به شمار می‌روند.

مطالعات دیگر نیز بر اهمیت استفاده از داده‌های چندمنبعی در بهبود پیش‌بینی عملکرد تأکید کرده‌اند. پژوهش‌های جابد و همکاران (۲۰۲۴)

sensitive, consumer-grade digital camera mounted on a commercial UAV to assess Bambara groundnut yield. *International Journal of Remote Sensing*. 2022; 43(2): 393–423. doi: 10.1080/01431161.2021.1974116

[7] Bassine FZ, Epule TE, Kechchour A, Chehbouni A. Recent applications of machine learning, remote sensing, and IoT approaches in yield prediction: A critical review. *arXiv [Preprint]*. 2023; arXiv:2306.04566. doi: 10.48550/arXiv.2306.04566

[8] Sharma RK, Kaur J, Feng G, Huang Y, Kumar C, Wang Y, et al. Maize and soybean yield prediction using machine learning methods: A systematic literature review. *Discover Agriculture*. 2025; 3(1): 64. doi: 10.1007/s44279-025-00215-6

[9] Jhajharia K, Mathur P, Jain S, Nijhawan S. Crop yield prediction using machine learning and deep learning techniques. *Procedia Computer Science*. 2023; 218: 406–17. doi: 10.1016/j.procs.2023.01.023

[10] Shawon SM, Ema FB, Mahi AK, Niha FL, Zubair H. Crop yield prediction using machine learning: An extensive and systematic literature review. *Smart Agricultural Technology*. 2025; 10:100718. doi: 10.1016/j.atech.2024.100718

[11] Oikonomidis A, Catal C, Kassahun A. Deep learning for crop yield prediction: A systematic literature review. *New Zealand Journal of Crop and Horticultural Science*. 2023; 51(1): 1–26. doi: 10.1080/01140671.2022.2032213

[12] Wang F, Li J, Zheng J, Chen S, Peng D, Zhang X, et al. Estimating soybean yields using causal inference and deep learning approaches with satellite remote sensing data. *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*. 2024; 17: 14161–14178. doi: 10.1109/JSTARS.2024.3435699

[13] Sabo F, Meroni M, Waldner F, Rembold F. Is deeper always better? Evaluating deep learning models for yield forecasting with small data. *Environmental Monitoring and Assessment*. 2023; 195(10): 1153. doi: 10.1007/s10661-023-11609-8

[14] Paudel D, de Wit A, Boogaard H, Marcos D, Osinga S, Athanasiadis IN. Interpretability of deep learning models for crop yield forecasting. *Computers and Electronics in Agriculture*. 2023; 206:107663. doi: 10.1016/j.compag.2023.107663

[15] Ashfaq M, Khan I, Afzal RF, Shah D, Ali S, Tahir M. Enhanced wheat yield prediction through integrated climate and satellite data using advanced AI techniques. *Scientific Reports*. 2025; 15(1): 18093. doi: 10.1038/s41598-025-02700-w

[16] Srivastava AK, Safaei N, Khaki S, Lopez G, Zeng W, Ewert F, et al. Winter wheat yield prediction using convolutional neural networks from environmental and phenological data. *Scientific Reports*. 2022; 12(1): 3215. doi: 10.1038/s41598-022-06249-w

[17] Jahan M. Optimization and validation of two machine learning algorithms for accurate prediction of irrigated wheat (*Triticum aestivum* L.) yield and identification of its influential

ابزاری کارآمد برای پیش‌بینی عملکرد محصول در مقیاس منطقه‌ای محسوب می‌شوند.

مشارکت نویسندگان

فاطمه بابایی به‌عنوان نویسنده اول، اجرای کامل مراحل تحلیل، مدل‌سازی و نگارش را بر عهده داشته است. دکتر علیرضا وفائی نژاد نظارت بر فرایند پژوهش و تأیید نهایی را انجام داده و دکتر علیرضا شریفی نیز در تأمین داده‌ها، ارزیابی نتایج و بازنگری ساختاری مقاله مشارکت فعال داشته‌اند. تمامی نویسندگان نسخه نهایی مقاله را مطالعه و تأیید کرده‌اند.

تشکر و قدردانی

این پژوهش با حمایت علمی و پژوهشی دانشگاه شهید بهشتی تهران، به‌ویژه گروه مهندسی نقشه‌برداری دانشکده مهندسی عمران، آب و محیط‌زیست، انجام شده است. مراتب سپاس و قدردانی خود را از تمامی اساتید ارجمندی که با ارائه راهنمایی‌های ارزشمند خود، نقش مؤثری در ارتقای کیفیت این مطالعه ایفا نمودند، ابراز می‌داریم.

تعارض منافع

«هیچ‌گونه تعارض منافع توسط نویسندگان بیان نشده است.»

منابع و مآخذ

- [1] Javed MA, Murad MAA. Crop yield prediction in agriculture: A comprehensive review of machine learning and deep learning approaches, with insights for future research and sustainability. *Heliyon*. 2024; 10(24):e40836. doi: 10.1016/j.heliyon.2024.e40836
- [2] Van Klompenburg T, Kassahun A, Catal C. Crop yield prediction using machine learning: A systematic literature review. *Computers and Electronics in Agriculture*. 2020; 177:105709. doi: 10.1016/j.compag.2020.105709
- [3] Parashar N, Johri P, Khan A, Gaur N, Kadry S. An integrated analysis of yield prediction models: A comprehensive review of advancements and challenges. *Computers, Materials & Continua*. 2024; 80(1): 389–425. doi: 10.32604/cmc.2024.050240
- [4] Bouras EH, Jarlan L, Er-Raki S, Balaghi R, Amazirh A, Richard B, et al. Cereal yield forecasting with satellite drought-based indices, weather data and regional climate indices using machine learning in Morocco. *Remote Sensing*. 2021; 13(16): 3101. doi: 10.3390/rs13163101
- [5] Okupska E, Juostas A, Gozdowski D, Wójcik-Gront E. Field-scale prediction of winter wheat yield using satellite-derived NDVI. *Agronomy*. 2026; 16(6): 670. doi: 10.3390/agronomy16060670
- [6] Jewan SYY, Pagay V, Billa L, Tyerman SD, Gautam D, Sparkes D, et al. The feasibility of using a low-cost near-infrared,

- [28] Soltaniyan M, Naderi Khorasgani M, Tadayon A. [Estimation of irrigated wheat (*Triticum aestivum* L.) yield using remote sensing data (case study: Shahrekord City)]. *Journal of Crop Production*. 2019; 12(2): 57–72. [In Persian] doi: 10.22069/ejcp.2019.14029.2066
- [29] Mohan RNJV, Pravalika Sree R, Sree RP. Next-gen agriculture: Integrating AI and XAI for precision crop yield predictions. *Frontiers in Plant Science*. 2025; 15:1451607. doi: 10.3389/fpls.2024.1451607
- [30] Baatz R. Generalization and feature attribution in machine learning models for crop yield and anomaly prediction in Germany. *arXiv [Preprint]*. 2025; arXiv:2512.15140. doi: 10.48550/arXiv.2512.15140
- [31] Guimarães N, Fraga H, Sousa JJ, Pádua L, Bento A, Couto P. Comparative evaluation of remote sensing platforms for almond yield prediction. *AgriEngineering*. 2024; 6(1): 240–258. doi: 10.3390/agriengineering6010015
- [32] Ang Y, Shafri HZM, Lee YP, Abidin H, Bakar SA, Hashim SJ, et al. A novel ensemble machine learning and time series approach for oil palm yield prediction using Landsat time series imagery based on NDVI. *Geocarto International*. 2022; 37(25): 9865–9896. doi: 10.1080/10106049.2022.2025920
- [33] Vannoppen A, Gobin A. Estimating yield from NDVI, weather data, and soil water depletion for sugar beet and potato in northern Belgium. *Water*. 2022; 14(8): 1188. doi: 10.3390/w14081188
- [34] Zhou Q, Ismaeel A. Integration of maximum crop response with machine learning regression model to timely estimate crop yield. *Geo-spatial Information Science*. 2021; 24(3): 474–483. doi: 10.1080/10095020.2021.1957723
- [35] Schils R, Olesen JE, Kersebaum K-C, Rijk B, Oberforster M, Kalyada V, et al. Cereal yield gaps across Europe. *European Journal of Agronomy*. 2018; 101: 109–120. doi: 10.1016/j.eja.2018.09.003
- [36] Kavipriya J, Vadivu G. Exploring crop yield prediction with remote sensing imagery and AI. In: 2024 International Conference on Advances in Computing, Communication and Applied Informatics (ACCAI); 2024. p. 1–5. doi: 10.1109/ACCAI61061.2024.10601854
- [37] Aslan MF, Sabanci K, Aslan B. Artificial intelligence techniques in crop yield estimation based on Sentinel-2 data: A comprehensive survey. *Sustainability*. 2024; 16(18): 8277. doi: 10.3390/su16188277
- factors in Khorasan Razavi Province. *Iranian Journal of Field Crops Research*. 2025; 23(4): 511–543. doi: 10.22067/jcresc.2025.93323.1395
- [18] Li X, Jin H, Eklundh L, Bouras EH, Olsson P-O, Cai Z, et al. Estimation of district-level spring barley yield in southern Sweden using multi-source satellite data and random forest approach. *International Journal of Applied Earth Observation and Geoinformation*. 2024; 134:104183. doi: 10.1016/j.jag.2024.104183
- [19] Sahu H, Garg PK, Vijay S, Dasgupta A. Predictive drivers and transferability of multi-scale machine learning based crop yield prediction under drought across European and Asian climates. *Agricultural Water Management*. 2026; 327:110254. doi: 10.1016/j.agwat.2026.110254
- [20] Mohammadi S, Rydgren K, Bakkestuen V, Gillespie MA. Impacts of recent climate change on crop yield can depend on local conditions in climatically diverse regions of Norway. *Scientific Reports*. 2023; 13(1): 3633. doi: 10.1038/s41598-023-30813-7
- [21] Liaghat A, Dehghani T, Rezaei Rad H, Ahmadpari H. [Estimating grain corn yield based on Landsat 8 satellite images (case study: Shahid Beheshti Agro-Industry lands of Dezful)]. *Iranian Journal of Irrigation and Drainage*. 2024; 18(3): 433–447. [In Persian]
- [22] Sardari Haroonieh A, Homaee M, Norouzi AA. [Evaluation of the FAO AquaCrop model for estimating wheat yield using Landsat satellite data]. *Iranian Journal of Irrigation and Drainage*. 2021; 15(5): 1212–1225. [In Persian]
- [23] Ghorbani M, Teymouri M, Salari Jazi M. [Wheat yield estimation using satellite images in Golestan Province]. *Agricultural Meteorology*. 2021; 9(1): 38–52. [In Persian] doi: 10.22125/agmj.2021.266547.1107
- [24] Amir Ashayeri E, Rezaverdinejad V, Bahmanesh J, Asadzadeh F, Rahimi M. [Modeling and forecasting rainfed wheat yield based on meteorological variables using combined artificial intelligence methods]. *Iranian Journal of Irrigation and Drainage*. 2025; 19(2): 243–261. [In Persian]
- [25] Taherihajivand A, Shirini K, Samadi Gharehveran S. [An overview of crop yield prediction using artificial intelligence algorithms]. *Journal of Agricultural Mechanization*. 2024; 9(3): 1–14. [In Persian] doi: 10.22034/jam.2024.61899.1276
- [26] Parviz L, Kazemi B, Hatef MA. [Utilizing climatic indices and a multi-criteria decision-making approach in crop yield prediction for agricultural policymaking]. *Journal of Drought and Climate Change Research*. 2024; 2(3): 49–66. [In Persian] doi: 10.22077/jdcr.2024.7942.1072
- [27] Ansari Ghoujghar M. [Development of a prediction toolbox for strategic wheat yield using machine learning algorithms to reduce food security risks (case study: Alborz Province)]. *Iranian Journal of Soil and Water Research*. 2022; 53(10): 2277–2294. [In Persian] doi: 10.22059/ijswr.2022.342638.669260

معرفی نویسندگان

AUTHOR(S) BIOSKETCHES

فاطمه بابایی دانشجوی کارشناسی ارشد مهندسی نقشه‌برداری گرایش سیستم‌های اداره زمین در دانشگاه شهید بهشتی است. علایق پژوهشی وی شامل سیستم‌های اطلاعات مکانی، هوش مکانی، یادگیری ماشین و کاربرد آن‌ها در کشاورزی هوشمند است. وی تاکنون یک مقاله

Water, and Environmental Engineering, Shahid Beheshti University, Tehran, Iran

✉ a_vafaei@sbu.ac.ir



علیرضا شریفی دانشیار گروه مهندسی نقشه‌برداری در دانشکده مهندسی عمران، آب و محیط‌زیست دانشگاه شهید بهشتی است. ایشان مدرک کارشناسی‌ارشد و دکتری خود را در رشته مهندسی نقشه‌برداری - سنجش‌ازدور به ترتیب در سال‌های ۱۳۸۷ و ۱۳۹۴ از

دانشگاه تهران دریافت کرده‌اند. فعالیت‌های پژوهشی اخیر وی در حوزه کاربرد یادگیری عمیق در پردازش تصویر فراطیفی و همچنین کاربردهای سنجش‌ازدور در پایش تغییرات اقلیمی، کشاورزی، مدیریت جنگل‌ها، مدیریت کاربری و پوشش اراضی و مدل‌سازی محیطی با استفاده از کلان‌داده‌های سنجش‌ازدور است. در حال حاضر ایشان بیش از ۱۰۰ مقاله علمی در مجلات منتشر کرده‌اند و عضو هیئت تحریریه نشریه IEEE TGRS است. همچنین ایشان در کمیته علمی و داوری بیش از ۴۰ مجله و کنفرانس علمی ملی و بین‌المللی فعالیت داشته‌اند. زمینه‌های تخصصی ایشان عبارت‌اند از: پردازش تصاویر فراطیفی و راداری، هوش مصنوعی در اطلاعات مکانی، توسعه کاربردهای سنجش‌ازدور.

Sharifi, A. Associate Professor at the Department of Geoinformation and Geomatics Engineering, Faculty of Civil, Water, and Environmental Engineering, Shahid Beheshti University, Tehran, Iran

✉ a_sharifii@sbu.ac.ir



علمی پژوهشی و چندین مقاله کنفرانسی ارائه کرده و همچنین سابقه دو ترم دستیار آموزشی و همیار پژوهشی را در کارنامه خود دارد.

Babayi, F. M.Sc. Student at the Department of Earth System Engineering, Faculty of Civil, Water and Environmental Engineering, Shahid Beheshti University, Tehran, Iran

✉ fa.babayi@mail.sbu.ac.ir



علیرضا وفائی نژاد دانشیار گروه مهندسی نقشه‌برداری در دانشکده مهندسی عمران، آب و محیط‌زیست دانشگاه شهید بهشتی است. وی دارای مدرک دکتری مهندسی عمران - نقشه‌برداری از دانشگاه صنعتی خواجه نصیرالدین طوسی است. تاکنون بیش از ۱۰۰

مقاله علمی در مجلات منتشر کرده‌اند. حوزه‌های پژوهشی اصلی ایشان شامل سیستم‌های اطلاعات مکانی، سیستم‌های اطلاعات مکانی هوشمند، یادگیری ماشین و کاربرد آن‌ها در مدیریت منابع آب و محیط‌زیست است. همچنین سوابق علمی و اجرایی متعددی از جمله مدیریت گروه مهندسی نقشه‌برداری و عضویت در هیئت تحریریه نشریات علمی در کارنامه ایشان وجود دارد.

Vafae-Nejad, A. Associate Professor at the Department of Geoinformation and Geomatics Engineering, Faculty of Civil,

Citation (Vancouver): Babayi F, Vafae-Nejad A, Sharifi A. [Final yield prediction of winter wheat using remote sensing data and machine-learning algorithms]. *J. RS. GEOINF. RES.* 2026; 4(1): 189-212

doi <https://doi.org/10.22061/jrsgr.2026.113038.1125>



COPYRIGHTS

© 2026 The Author(s). This is an open-access article distributed under the terms and conditions of the Creative Attribution-NonCommercial 4.0 International (CC BY-NC 4.0)
(<https://creativecommons.org/licenses/by-nc/4.0/>)