



Crop Classification and Acreage Estimation at Basin Scale Using Different Machine Learning Approaches

Masoud Soltani¹ | Bahareh Bahmanabadi¹ | Ali Mokhtaran²

¹ Department of Water Science and Engineering, Faculty of Agriculture and Natural Resources, Imam Khomeini International University, Qazvin, Iran. Email: msoltani@eng.ikiu.ac.ir

*1Corresponding Author: Bahareh Bahmanabadi, Department of Water Science and Engineering, Faculty of Agriculture and Natural Resources, Imam Khomeini International University, Qazvin, Iran. Email: b.bahmanabadi@gmail.com

² Agricultural Engineering Research Department, Khuzestan Agricultural and Natural Resources Research and Education Center (AREEO), Ahwaz, Iran. alimokhtaran@gmail.com

Article Info

Article type: Research Article

Article history:

Received: Sep. 22, 2025

Revised: Nov. 15, 2025

Accepted: Dec. 13, 2025

Published online: Jan. 2026

Keywords:

**Marun-Jarrahi Basin,
Random Forest,
Sentinel,
Support Vector Machine,
XGBoost**

ABSTRACT

Accurate mapping of crop types and their spatial distribution, along with estimating cultivated areas, plays a critical role in water resource planning and agricultural land management in arid and semi-arid regions. In this study, field surveys were first conducted to collect precise geographic coordinates of selected crop fields, including wheat, maize, sesame, and alfalfa, and these reference points were subsequently integrated into the Google Earth Engine environment for alignment with remote sensing data. Sentinel-2 imagery was processed through atmospheric correction and cloud masking, while Sentinel-1 data underwent geometric correction and speckle filtering. A set of spectral indices derived from Sentinel-2, together with backscatter intensity and polarization ratio features extracted from Sentinel-1, was generated and integrated to enhance the separability of crop types with similar phenological characteristics. The collected reference samples were divided into training (70%) and validation (30%) subsets, and three classification algorithms, Random Forest, Support Vector Machine, and Extreme Gradient Boosting, were applied to produce the final crop type map. Model performance was evaluated using the confusion matrix, overall accuracy, and the Kappa coefficient. The results indicated that the Random Forest algorithm achieved the highest performance, with an overall accuracy of 96% and a Kappa coefficient of 0.93, demonstrating superior discrimination of crop types, particularly wheat. The Extreme Gradient Boosting model ranked second with an accuracy of 89%, while the Support Vector Machine exhibited lower performance. Comparison of the estimated cultivated areas with official agricultural statistics further confirmed the high reliability of the Random Forest results.

Cite this article: Soltani, M., Bahmanabadi, B., Mokhtaran, A., (2026) Crop Classification and Acreage Estimation at Basin Scale Using Different Machine Learning Approaches, *Iranian Journal of Soil and Water Research*, 56 (11), 3129-3156. <https://doi.org/10.22059/ijswr.2025.402830.670012>

© The Author(s).

Publisher: University of Tehran Press.

DOI: <https://doi.org/10.22059/ijswr.2025.402830.670012>



EXTENDED ABSTRACT

Introduction

Accurate mapping of crop types and the estimation of their cultivated areas are essential components of sustainable agricultural management, particularly in arid and semi-arid regions such as southwestern Iran. Rapid population growth, diminishing water resources, and the accelerating impacts of climate change have intensified the need for timely and reliable information on cropping patterns. Remote sensing provides a practical and cost-efficient means for agricultural monitoring, offering repeated temporal coverage and high spatial detail through multisource satellite observations. The combined use of optical and radar imagery has proven especially valuable, as it integrates both spectral and structural information to improve discrimination among crops with similar spectral characteristics.

Although considerable progress has been made in crop mapping using satellite imagery, there remains a lack of consistent and transferable frameworks that combine optical and radar data in an operational manner at basin scale in Iran. The Maroon–Jarrahi Basin, one of the country's key agricultural regions, faces severe challenges from groundwater depletion and increasing climatic variability. This situation requires an accurate, spatially explicit understanding of crop distribution, particularly for major crops such as wheat, maize, sesame, and alfalfa. The present study seeks to establish a comprehensive classification framework that integrates Sentinel-1 and Sentinel-2 datasets to generate accurate crop type maps and reliable acreage estimates for this critical agricultural region.

Methods

The study employed a supervised classification design using multisource satellite data and extensive field observations collected between 2021 and 2024 within the Maroon–Jarrahi sub-basin of Khuzestan Province, southwestern Iran. Field data included precise GPS-located samples of four major crops (wheat, maize, sesame, and alfalfa) collected across several administrative districts.

Sentinel-2 imagery was atmospherically corrected and cloud-masked using the Fmask algorithm. A suite of vegetation indices sensitive to plant biophysical and phenological traits was derived and resampled to a 10-m spatial resolution. For Sentinel-1, radar backscatter data acquired in Interferometric Wide (IW) mode with VV and VH polarizations were processed. Speckle noise was minimized using a 5×5 adaptive Lee filter, preserving textural integrity. The polarization ratio (VV/VH) was further computed to capture variations in canopy structure and surface moisture. All radar and optical layers were fused into a single composite dataset to maximize spectral and structural feature diversity. Moreover, All samples for each class were compiled, with 70% allocated for model training and 30% reserved for independent validation.

Classification experiments were conducted in the Google Earth Engine environment coupled with Python-based processing. Three classification algorithms, Random Forest (RF), Support Vector Machine (SVM), and Extreme Gradient Boosting (XGBoost) were calibrated and tested. For RF, model parameters included 100 decision trees, five variables per split, and a minimum leaf population of three. The SVM classifier used a radial basis kernel, with optimized γ values ranging from 0.0002 to 0.0004 and C values between 1 and 100. The XGBoost model underwent grid-based hyperparameter tuning with five-fold cross-validation to identify optimal values of $n_estimators$, max_depth , $learning_rate$, and $subsample_rate$.

Model evaluation was based on a confusion matrix and a set of standard accuracy metrics, including overall accuracy (OA), Kappa coefficient, precision, recall, and F1-score. A spatially stratified random sampling approach was adopted to minimize spatial autocorrelation between training and validation samples. The final classified maps were exported as GeoTIFF files, and crop areas were computed at both county and basin levels.

Results

The comparative assessment of classifiers revealed clear differences in performance. The Random Forest model produced the highest accuracy (OA = 96%, Kappa = 0.93, F1 = 0.97), demonstrating a strong ability to distinguish spectrally similar crops, especially between wheat and alfalfa, while maintaining resilience to noise and spatial heterogeneity. The RF approach effectively reduced confusion between vegetated and non-vegetated classes, with minimal misclassification of water bodies or algae.

The XGBoost classifier ranked second (OA = 89%, Kappa = 0.83, F1 $\approx$ 0.90), achieving stable results across most crop classes but showing some confusion between bare soil and sparse vegetation, particularly under post-harvest conditions when spectral signatures converge. The Support Vector Machine yielded the lowest performance (OA $\approx$ 54%, Kappa = 0.31, F1 $\approx$ 0.55), largely due to its sensitivity to inter-class spectral overlap, notably between sesame and maize.

Area estimation results confirmed strong agreement between RF-derived crop acreage and official

agricultural statistics, with deviations of less than 5% for wheat. For alfalfa, a minor overestimation was observed, attributed to its persistent greenness and mixed-pixel effects. The combined use of Sentinel-1 and Sentinel-2 data enhanced the accuracy of water-related classifications, accurately delineating open water and aquatic vegetation. Visual assessment further showed that the RF-generated maps exhibited the most coherent and spatially consistent patterns, whereas XGBoost maps displayed moderate fragmentation and SVM outputs contained irregular artifacts. Overall, the integration of radar and optical data markedly improved class separability, reinforcing the complementary value of multisensor fusion for crop mapping in arid environments.

Conclusion

This research establishes a coherent and replicable framework for crop classification and acreage estimation based on the integration of optical and radar satellite data. The results highlight the potential of Sentinel-1 and Sentinel-2 fusion to significantly improve the accuracy and reliability of crop mapping in data-scarce arid and semi-arid regions. Among the tested classification approaches, the Random Forest algorithm provided the most consistent and robust outcomes, producing highly accurate and spatially stable crop maps.

The proposed methodology provides a transferable and scalable foundation for regional agricultural monitoring, enabling timely assessment of cropping patterns and supporting more informed water resource management and agricultural policy decisions. Future efforts should incorporate multi-year climatic and phenological datasets to strengthen temporal generalization and move toward dynamic, near-real-time monitoring systems across diverse agro-ecological zones of Iran.

Author Contributions

Masoud Soltani: Guiding, Conceptualization, Editing the paper, Controlling the results

Bahareh Bahmanabadi: Programming, Conceptualization, methodology, Writing - Initial draft preparation, performing software/statistical analysis

Ali Mokhtaran: supplying the field data in the Marun Basin, Guiding, Conceptualization

Data Availability Statement

Data available on request from the authors.

Acknowledgements

The authors would like to express their sincere gratitude to Imam Khomeini International University for providing financial and logistical support that greatly contributed to the success of this research. We also acknowledge the valuable assistance of the Agricultural Engineering Research Institute (AERI) for supplying the field data in the Marun Basin.

Ethical considerations

The authors avoided data fabrication, falsification, plagiarism, and misconduct. Conflict of interest: The author declares no conflict of interest.

Conflict of interest

The authors of this article declared no conflict of interest regarding the authorship or publication of this article.

طبقه‌بندی محصولات کشاورزی و برآورد سطح کشت در مقیاس حوضه‌ای با رویکرد یادگیری ماشین

مسعود سلطانی^۱ | بهاره بهمن‌آبادی^۲ | علی مختاران^۳

۱. گروه علوم و مهندسی آب، دانشکده کشاورزی و منابع طبیعی، دانشگاه بین‌المللی امام خمینی (ره)، قزوین، ایران. رایانامه:

msoltani@eng.ikiu.ac.ir

۲. نویسنده مسئول، گروه علوم و مهندسی آب، دانشکده کشاورزی و منابع طبیعی، دانشگاه بین‌المللی امام خمینی (ره)، قزوین، ایران.

رایانامه: b.bahmanabadi@alumni.ut.ac.ir

۳. بخش تحقیقات فنی و مهندسی کشاورزی، مرکز تحقیقات و آموزش کشاورزی و منابع طبیعی خوزستان، سازمان تحقیقات، آموزش و ترویج

کشاورزی، اهواز، ایران رایانامه: alimokhtaran@gmail.com

اطلاعات مقاله	چکیده
نوع مقاله: مقاله پژوهشی	در مناطق خشک و نیمه‌خشک، تهیه نقشه دقیق نوع و پراکندگی محصولات زراعی و برآورد سطح زیرکشت آن‌ها نقشی اساسی در برنامه‌ریزی منابع آب و مدیریت بهینه اراضی دارد. در این پژوهش، ابتدا پایش‌های میدانی در حوضه مارون-جراحی، انجام و مختصات مزارع محصولات منتخب شامل گندم، ذرت، کنجد و یونجه به‌صورت دقیق برداشت شد و سپس این نقاط مرجع در محیط گوگل ارث انجین به‌منظور همسان‌سازی با داده‌های سنجش‌ازدور مورد استفاده قرار گرفتند. در ادامه، تصاویر سنتینل-۲ پس از تصحیح جوی و حذف ابر و تصاویر سنتینل-۱ پس از اعمال فیلتر اسپکل و تصحیح هندسی آماده شدند. مجموعه‌ای از شاخص‌های طیفی برگرفته از سنتینل-۲ و ویژگی‌های شدت و نسبت پلاریزاسیون برگرفته از سنتینل-۱ استخراج و با یکدیگر ترکیب گردید تا امکان تفکیک بهتر محصولات با رفتار فنولوژیک مشابه فراهم شود. نمونه‌های میدانی به نسبت ۷۰ درصد برای آموزش و ۳۰ درصد برای ارزیابی مدل‌ها تقسیم شدند و سه الگوریتم جنگل تصادفی، SVM و XGBoost برای طبقه‌بندی نهایی به‌کار گرفته شد. ارزیابی مدل‌ها بر اساس ماتریس درهم‌ریختگی، صحت کلی و ضریب کاپا انجام شد. نتایج نشان داد الگوریتم جنگل تصادفی با صحت کلی ۹۶٪ و ضریب کاپا ۰/۹۳ و F1-Score، ۰/۹۷ عملکرد برتری نسبت به دو روش دیگر داشته و تفکیک گندم و سایر محصولات را با دقت بالا ممکن ساخته است؛ در حالی که مدل XGBoost با صحت ۸۹٪ در رتبه بعدی قرار گرفت و SVM عملکرد ضعیف‌تری نشان داد. مقایسه سطح زیرکشت برآورد شده با آمار رسمی نیز تطابق قابل توجه نتایج جنگل تصادفی را تأیید کرد.
واژه‌های کلیدی: حوضه مارون-جراحی، سنتینل، جنگل تصادفی، svm XGBoost	

استناد: سلطانی؛ مسعود، بهمن‌آبادی؛ بهاره، مختاران؛ علی، (۱۴۰۴) طبقه‌بندی محصولات کشاورزی و برآورد سطح کشت در مقیاس حوضه‌ای با رویکرد یادگیری ماشین،

مجله تحقیقات آب و خاک ایران، ۵۶ (۱۱)، ۳۱۲۹-۳۱۵۶. <https://doi.org/10.22059/ijswr.2025.402830.670012>

© نویسندگان.

ناشر: مؤسسه انتشارات دانشگاه تهران.

DOI: <https://doi.org/10.22059/ijswr.2025.402830.670012>

مقدمه

رشد جمعیت و تشدید اثرات تغییر اقلیمی در سال‌های اخیر، اهمیت مدیریت و پایش در کشاورزی و منابع آب را دوچندان کرده‌اند. به‌منظور سازگاری با تغییرات اقلیمی، تضمین امنیت غذایی، تدوین سیاست‌های کشاورزی و مطالعات اکولوژیکی کشاورزی، دانستن توزیع مکانی و زمانی کشت محصولات، اطلاعات ضروری و پایه‌ای در اختیار ما قرار می‌دهد (Gessner et al., 2025). بنابراین تولید نقشه‌های دقیق و به‌روز الگوی کشت، نقشی کلیدی در مدیریت و برنامه‌ریزی کشاورزی پایدار ایفا می‌کند. در دهه‌های اخیر سنجش از دور پیشرفت‌های بسیاری در زمینه‌های مختلف داشته‌است. با ظهور ماهواره‌های دارای سنجنده با قدرت تفکیک زمانی (روزانه تا هر ۱۶ روز یکبار) و کیفیت مکانی (کمتر از ۱۰ متر) اطلاعات دقیقی را ارائه می‌دهند. همچنین، با ارتقای کیفیت طیفی و افزایش تعداد باندها، دید وسیع‌تری از پوشش سطح زمین و تغییرات کاربری اراضی در اختیار محققان قرار می‌دهد (Reisi-Gahrouei et al., 2019). بنابراین، سنجش از دور ابزاری کاربردی و مقرون به‌صرفه برای مطالعه و تحلیل موضوعات مختلف در کشاورزی و محیط‌زیست تبدیل شده‌است. علاوه بر این، رویکرد تلفیق سنجنده‌های مختلف از جمله اپتیک و راداری انجام تحلیل‌های ترکیبی را فراهم کرده که دقت و کارایی نتایج را به‌طور چشمگیری افزایش می‌دهد (Reisi Gahrouei et al., 2019). بعد از سال ۲۰۱۷، داده‌های ماهواره‌های سنتینل-۱ و سنتینل-۲ به‌طور همزمان در دسترس قرار گرفت و این امر منجر به تولید سری‌های زمانی متراکم با تناوب پنج‌روزه و دقت مکانی ۱۰ تا ۲۰ متر شد. چنین داده‌هایی کمک شایانی به نقشه‌برداری دقیق اراضی کشاورزی کرده‌است؛ چراکه پیوستگی زمانی بالای تصاویر، تغییرات زمانی و روند رشد محصولات را به‌خوبی نشان می‌دهد. این ویژگی نقش کلیدی در تشخیص و تفکیک انواع محصولات دارد و علاوه بر آن، داده‌ها در مقیاس مزرعه ارائه می‌شوند که به افزایش دقت و کارایی تحلیل‌های کشاورزی کمک شایانی می‌کند (Phiri et al., 2020).

با توجه به اهمیت تولید نقشه‌های به‌روز الگوی کشت، رویکردهای مختلفی در سال‌های اخیر توسعه یافته‌اند. یکی از نخستین مدل‌های توسعه‌یافته برای شناسایی محصولات کشاورزی، بر پایه‌ی دوره رشد گیاه و ویژگی‌های طیفی متمایزکننده‌ی آن، از جمله میزان کلروفیل و الگوی بازتاب در باندهای مختلف طیفی، استوار است. مطالعات متعددی نشان داده‌اند که تغییرات فیزیولوژیکی و طیفی گیاه در طول دوره رشد، می‌تواند به‌طور مؤثری موجب تفکیک آن از سایر محصولات گردد (Ashourloo et al., 2022; Dong et al., 2020; Liang et al., 2024; Haddai, et al. 2025). یکی از چالش‌های اصلی در روش‌های بررسی دوره رشد گیاه، نیاز به داده‌های زمینی دقیق برای آموزش و ارزیابی مدل‌ها است. این در حالی است که مدل‌های یادگیری ماشین می‌توانند با استفاده از داده‌های سنجش از دوری و نیاز کمتر به داده‌های زمینی متنوع و فراوان، عملکرد قابل قبولی داشته باشند، به‌ویژه در مناطق با دسترسی محدود به داده‌های زمینی می‌تواند بسیار کمک‌کننده باشد (Halder et al., 2025). از سویی دیگر، استفاده از مدل‌های یادگیری ماشین، تحلیل مجموعه داده‌های بزرگ مقیاس و پیچیده را بهبود بخشیده و این امر به پیش‌بینی دقیق‌تر در مورد بهره‌وری محصولات و برآورد نیاز آبی کمک کرده است (Sarker, 2021). در مدل‌های یادگیری ماشین، قابلیت یادگیری از داده‌ها، آن‌ها را به ابزاری مناسب برای طبقه‌بندی تصاویر تبدیل کرده و راه‌حلی مناسب برای کشاورزی هوشمند ارائه می‌دهد. در این بین، روش‌های نظارت‌شده، که از داده‌های چندزمانه و شاخص‌های پوشش گیاهی استفاده می‌کنند، نتایج قابل‌توجهی را در تمایز انواع محصولات کشاورزی ارائه کرده‌اند (Yadav et al., 2024). از جمله تحقیقات صورت گرفته در این زمینه می‌توان به تحقیق (Şimşek, 2025) اشاره کرد، در این مطالعه، مقایسه‌ای از عملکرد مدل‌های RF، شبکه عصبی مصنوعی (ANN)، SVM و XGBoost در طبقه‌بندی چندزمانه محصولات کشاورزی با استفاده از تصاویر سنتینل-۲ در ناحیه چوکورووا واقع در ترکیه، انجام شده‌است. نتایج نشان داد که مدل XGBoost با صحت کلی ۹۲/۱۴٪ بهترین عملکرد در مقایسه با سایر مدل‌ها را داشته است (Şimşek, 2025). همچنین، Zheng و همکاران، (۲۰۲۴) از یک رویکرد ترکیبی برای نقشه‌برداری محصولات کشاورزی در سامانه خرده مالکی مناطق کوهستانی در جنوب غربی چین استفاده کردند. این روش با ادغام تصاویر لندست ۸، سنتینل-۱ و سنتینل-۲ و بهره‌گیری از شبکه عصبی عمیق CNN-2D، توانست با صحت کلی ۹۹/۹۳٪، ضریب کاپا ۰/۹۲۵۳ و F1-score بالا، عملکرد بهتری نسبت به روش‌های RF، SVM و XGBoost ارائه دهد. این نتایج نشان‌دهنده پتانسیل بالای رویکردهای ترکیبی تصاویر ماهواره‌ای برای طبقه‌بندی دقیق محصولات در مناطق کوهستانی و شرایط طبیعی دشوار است. در تحقیقی دیگر، عملکرد

1 SAR

2 Random Forest

3 Support Vector Machine

4 eXtreme Gradient Boosting



الگوریتم‌های یادگیری ماشین برای طبقه‌بندی پوشش گیاهی تالاب‌های مختلف در اروپا با استفاده از داده‌های چندزمانه سنتینل-۲ ارزیابی شد. نتایج نشان داد که روش‌های تجمیعی مانند RF و XGBoost عموماً قدرت پیش‌بینی بالاتری دارند، در حالی که SVM بیشترین دقت کلی و F-score را برای تمامی کلاس‌ها کسب کرد (Piaser & Villa, 2023). Zaka and Samat, 2024. نشان دادند که استفاده از داده‌های سنجش از دور ترکیبی همراه با مدل‌های یادگیری ماشین مانند SVM و RF، دقت و کارایی بالایی در شناسایی و نقشه‌برداری محصولات ارائه می‌دهد (Zaka & Samat, 2024).

در ایران، ذرتی‌پور و همکاران ۱۴۰۴، به بررسی و مدل‌سازی تغییرات پوشش گیاهی جنگل‌های جنوب غرب ایران با استفاده از تصاویر ماهواره‌ای لندست ۸ و با استفاده از شش الگوریتم یادگیری ماشین شامل جنگل تصادفی، بردار پشتیبان، مدل جمعی تعمیم یافته، مدل خطی تعمیم یافته، درخت طبقه بندی و رگرسیون و رگرسیون تقویت شده درختی پرداختند، نتایج نشان داد که مدل جنگل تصادفی دارای صحت و دقت بیشتری نسبت به پنج مدل دیگر بود (ZoratiPour et al., 2025). سلطانی و بهمن آبادی ۱۴۰۴، مطالعه‌ای با هدف ارزیابی دقت طبقه‌بندی تصاویر ماهواره‌ای سنتینل-۱ و ۲ برای دوره‌های کشت بهاره و پاییزه در سال ۱۴۰۲-۱۴۰۳ براساس سه الگوریتم یادگیری ماشین شامل جنگل تصادفی، بردار پشتیبان و گرادیان تقویت شده در منطقه دشت قزوین انجام شده است. الگوریتم تقویت گرادیان شدید نیز با صحت کلی ۹۳/۹۴ درصد عملکرد قابل قبولی داشت، اما برای کلاس‌هایی با شباهت طیفی، دقت کمتری نسبت به جنگل تصادفی نشان داد. در مقابل، روش ماشین بردار پشتیبان به دلیل حساسیت به هم‌پوشانی طیفی، در جداسازی کلاس‌ها عملکرد ضعیف‌تری داشت و در نهایت روش طبقه‌بندی جنگل تصادفی با دقت کلی و ضریب کاپای بیش از ۰/۹۹ بهترین عملکرد را داشت (سلطانی، م. و بهمن‌آبادی، ب.، ۱۴۰۴). خسروی، ۱۴۰۳، با استفاده از سری زمانی تصاویر ماهواره لندست-۸ و الگوریتم‌های یادگیری ماشین پیشرفته یک چهارچوب تهیه نقشه نوع محصول کشاورزی مرودشت استان فارس ارائه داد نتایج نشان داد که روش‌های آنالیز انحراف زمانی پویا و جنگل تصادفی نسبت به روش‌های دیگر کارایی بسیار بیشتری (با افزایش صحت کلی به میزان ۱۰٪ تا ۱۲٪ بیشتر) در تهیه نقشه نوع محصول کشاورزی منطقه مطالعه شده داشتند (خسروی، ۱۴۰۳). باقری و همکاران، ۱۴۰۳ به پیش سطح زیر کشت و طبقه بندی محصولات زراعی با استفاده از فناوری سنجش از دور ماهواره ای در منطقه سیلاخور پایین پرداختند. براساس نتایج، مساحت اراضی گندم و جو منطقه، ۱۰۶۳۹ هکتار برآورد شد که در مقایسه با آمارنامه جهادکشاورزی یعنی ۵۶/۹۹۶ هکتار، خطایی حدود ۸/۶٪ را نشان می‌دهد. برای شناسایی و تهیه نقشه کشت محصولات کشاورزی مختلف از دو الگوریتم ماشین بردار پشتیبان و شبکه‌عصبی مصنوعی استفاده شد و محصولات در پنج طبقه، کشت تابستانه (صیفی جات: خیار، گوجه، هندوانه، خربزه)، مزارع کشت غلات زمستانه (گندم و جو دیم)، مزارع کشت آبی (یونجه، گندم، جو و کلزای آبی، سیب زمینی)، کشت برنج و اراضی غیر کشاورزی طبقه‌بندی شدند. کمترین مقدار خطای کمیسیون در هر دو روش طبقه‌بندی، ۱/۸٪، مربوط به کشت تابستانه است و خطای امیسیون در اراضی غیر کشاورزی در روش ماشین بردار پشتیبان ۵ و در روش شبکه عصبی مصنوعی ۵/۲٪ بود (باقری و همکاران، ۱۴۰۳).

در پژوهش نویدی و همکاران، ۱۴۰۰ از تصاویر چند زمانه سنتینل-۲ در تفکیک و تعیین سطح زیر کشت اراضی کشاورزی منطقه بسطام با کمک دوره فنولوژیکی استفاده شد. براساس نتایج بدست آمده، مدل ماشین بردار پشتیبان برای گندم بیش‌ترین مساحت (۳۴۲۳ هکتار) و برای ذرت علوفه‌ای کمترین مساحت (۷۳۸ هکتار) به‌دست آمد. نتایج حاصله نشان داد که استفاده از تصاویر ماهواره‌ای چند زمانه و شاخص فنولوژیکی، توانایی قابل قبولی را در تفکیک محصولات زراعی و تعیین مساحت و تهیه نقشه اراضی کشاورزی داشت (نویدی و همکاران، ۱۴۰۰). طاهری دهکردی و ولدان زوج، ۱۴۰۰ به طبقه‌بندی کاربری اراضی کشاورزی با استفاده از تصاویر ماهواره‌ای سنتینل و شبکه عصبی پیچشی سه‌بعدی عمیق (مطالعه موردی: شهرکرد) پرداختند، روش ارائه شده با دو نوع زمانی و زمانی-مکانی روش‌های جنگل تصادفی و ماشین بردار پشتیبان نیز مورد مقایسه قرار گرفت که با اختلاف حداقل ۲/۴ درصدی، عملکرد بهتری از خود نشان داد (طاهری دهکردی، ع. و ولدان زوج، م.ج. ۱۴۰۰). در تحقیق قدسی و همکاران ۱۳۹۹، به مقایسه دقت روش‌های ماشین بردار پشتیبان و جنگل تصادفی در تهیه نقشه کاربری اراضی و محصولات زراعی، با استفاده از تصاویر چندزمانه سنتینل-۲ در منطقه دشت سنجابی روانسر پرداختند، در این تحقیق، ز تصاویر چندزمانه سنتینل-۲ و روش‌های طبقه‌بندی ماشین بردار پشتیبان و جنگل تصادفی استفاده و دقت آنها با یکدیگر مقایسه شد. ارزیابی دقت‌ها نشان داد ماشین بردار پشتیبان، با صحت کلی ۹۱/۳۶٪ و ضریب کاپای ۰/۸۹۲۷، نقشه کاربری اراضی و محصولات دقیق‌تری، در قیاس با روش جنگل تصادفی، تولید می‌کند (قدسی و همکاران، ۱۳۹۹).

اگرچه تاکنون مطالعات متعددی با هدف طبقه‌بندی محصولات کشاورزی از داده‌های ماهواره‌ای سنتینل بهره گرفته‌اند، اما در حوضه آبریز مارون-جراحی تاکنون چارچوبی جامع و یکپارچه ارائه نشده است که بتواند داده‌های راداری سنتینل-۱ و اپتیکی سنتینل-۲ را در یک

بازه زمانی به‌صورت هم‌زمان تلفیق نماید و از مجموعه‌ای از شاخص‌های طیفی و راداری منتخب برای تفکیک محصولات با شباهت طیفی بالا استفاده کند. در این منطقه همچنین پژوهشی گزارش نشده است که بر پایه نمونه‌برداری میدانی گسترده و دقیق، مدل‌های یادگیری ماشین را آموزش داده و دقت طبقه‌بندی و برآورد سطح زیرکشت را به‌طور هم‌زمان با داده‌های رسمی مقایسه کند.

استان خوزستان و به‌ویژه زیرحوضه آبریز مارون-جراحی، یکی از مهم‌ترین قطب‌های کشاورزی کشور به‌شمار می‌رود و نقش راهبردی در تولید محصولات کلیدی نظیر گندم، ذرت و خرما دارد. دسترسی به اطلاعات دقیق و به‌روز درباره سطح زیرکشت و الگوی کشت در این منطقه، نقشی اساسی در برنامه‌ریزی کشاورزی، مدیریت پایدار منابع آب و ارتقای امنیت غذایی کشور ایفا می‌کند. با توجه به چالش‌های فزاینده‌ای نظیر کاهش منابع آبی، نوسانات بارندگی و اثرات محسوس تغییر اقلیم (Ghanian, 2024; Hajirad et al., 2023; Raeisi et al., 2022)، توسعه روش‌های دقیق و کارآمد برای پایش سطح زیرکشت در این منطقه ضرورت دارد. این داده‌های پایه، زیربنای تصمیم‌سازی در حوزه‌هایی همچون تخصیص منابع آبی، پایش بین‌تولید، تنظیم بازار محصولات کشاورزی و پایش اثرات خشکسالی به‌شمار می‌روند. (Foley et al., 2011)

بر این اساس، پژوهش حاضر با هدف ارزیابی و مقایسه سه الگوریتم پرکاربرد یادگیری ماشین شامل ماشین بردار پشتیبان (SVM)، جنگل تصادفی (Random Forest) و تقویت گرادیان شدید (XGBoost) در طبقه‌بندی محصولات زراعی حوضه مارون-جراحی انجام گرفته است.

نوآوری اصلی این پژوهش، علاوه بر کاربرد مدل‌های پیشرفته یادگیری ماشین در یکی از مهم‌ترین قطب‌های کشاورزی ایران، در ارائه یک چارچوب یکپارچه برای تلفیق ویژگی‌های اپتیکی و راداری به منظور تفکیک محصولات با شباهت طیفی بالا در شرایط اقلیمی منطقه است. این رویکرد علاوه بر افزایش دقت طبقه‌بندی، مسیر تازه‌ای برای توسعه مدل‌های ترکیبی در پایش پویای الگوی کشت، کشاورزی دقیق و مدیریت منابع آب در ایران و مناطق خشک مشابه ارائه می‌کند.

در ادامه مقاله، ابتدا محدوده مطالعه و داده‌های مورد استفاده معرفی می‌شود؛ سپس روش‌شناسی تحقیق شامل مراحل پیش‌پردازش، استخراج ویژگی‌ها و پیاده‌سازی مدل‌های یادگیری ماشین تشریح می‌گردد. در بخش بعد، نتایج حاصل از ارزیابی مدل‌ها و تحلیل دقت طبقه‌بندی ارائه می‌شود و در پایان، پیشنهادهایی برای جهت‌گیری پژوهش‌های آتی مطرح خواهد شد.

روش‌شناسی

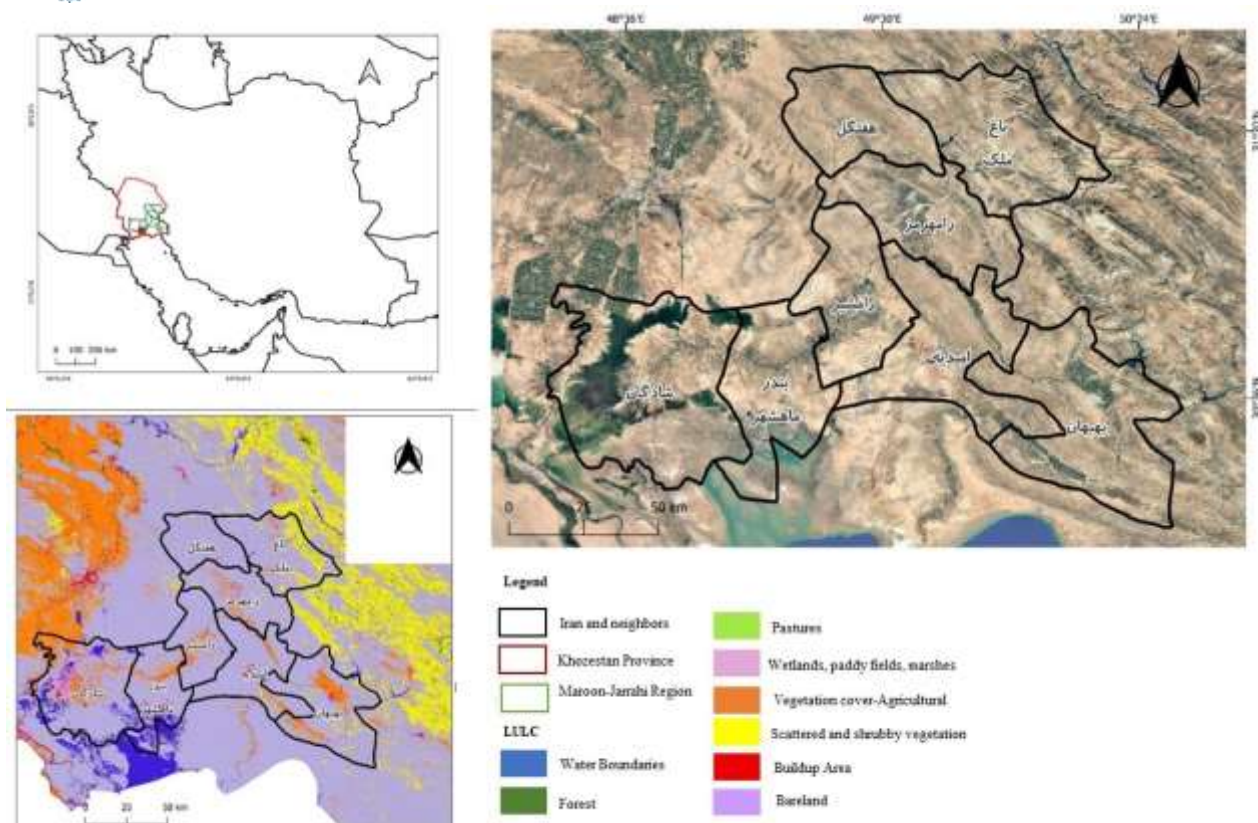
منطقه مورد مطالعه

زیرحوضه مارون - جراحی خوزستان در جنوب غرب ایران قرار دارد و در محدوده جغرافیایی $50^{\circ}05'$ الی $51^{\circ}11'$ طول شرقی و $30^{\circ}39'$ الی $31^{\circ}21'$ عرض شمالی واقع شده است. رودخانه مارون توسط سرشاخه‌های سقاوله، لوداب، شور و چاروساق از ارتفاعات زاگرس سرچشمه گرفته و پس از طی مسیری طولانی از طریق تنگ تکاب وارد دشت بهبهان می‌شود (Zalaki Badili et al., 2013). محدوده مورد مطالعه زیرحوضه آبخیز رودخانه مارون دارای مساحتی حدود 2750 کیلومتر مربع و حداکثر ارتفاع در این حوضه 3483 متر و حداقل ارتفاع 364 متر در نزدیکی سد مارون می‌باشد. اهمیت کیفیت این رودخانه از آنجا مشخص می‌شود که یکی از سرشاخه‌های رود جراحی است که در مسیر خود آب مورد نیاز شهرها و روستاهای زیاد و نیز هزاران هکتار اراضی کشاورزی، باغ‌ها، نخلستان‌ها و کارخانه‌های صنعتی را تأمین می‌کند و در نهایت به تالاب شادگان منتهی می‌گردد (شریفی پیچون، م. و همکاران، ۱۳۹۸).

طبق گزارش سازمان هواشناسی کشوری، منطقه مورد بررسی براساس روش دومارتن دارای اقلیم خشک تا نیمه‌خشک بوده و تنوع بالایی از محصولات کشاورزی و نخیلات را دربرمی‌گیرد (شکل ۱). در جدول ۱، اطلاعات محصولات مورد مطالعه در این منطقه شامل، گندم، یونجه، ذرت و کنجد ارائه شده است.

جدول ۱. تقویم زراعی محصولات مورد بررسی و نوع کشت

محصول	تاریخ کاشت (به صورت تقریبی در منطقه)	برداشت	نوع کشت
ذرت	۲۰ تیر تا ۱۵ مرداد	۱۵ آبان تا اوایل آذر	آبی
کنجد	اول تیر تا ۱۵ تیر	۱۵ شهریور تا اوایل مهر	آبی
یونجه	اوایل مهر	سالانه با حدود ۱۰ چین در سال	آبی
گندم	اواخر مهر تا اواسط آبان	۱ تا ۱۵ اردیبهشت	دیم و آبی



شکل ۱. کاربری اراضی و موقعیت جغرافیایی منطقه مطالعاتی

تصاویر ماهواره

در این مطالعه، به منظور تهیه نقشه‌های طبقه‌بندی محصولات کشاورزی، از ترکیب داده‌های اکتیکی سنتینل-۲ و داده‌های راداری سنتینل-۱ در بازه زمانی ۱۴۰۰-۱۴۰۳ استفاده شد. تمامی مراحل دسترسی به آرشیو تصاویر، پیش‌پردازش و استخراج ویژگی‌ها با استفاده از پلتفرم پردازش ابری گوگل ارث انجین انجام پذیرفت. مجموعه ماهواره‌های سنتینل که نخستین عضو آن در سال ۲۰۱۴ پرتاب گردید، به دلیل قدرت تفکیک مکانی مناسب، تناوب تصویربرداری تقریباً هر ۵-۶ روز (با استفاده همزمان از سکوها A/B و دسترسی رایگان به داده‌ها، کاربرد گسترده‌ای در مطالعات کشاورزی، پایش پوشش گیاهی و تغییرات کاربری اراضی در سطح جهانی یافته‌اند.

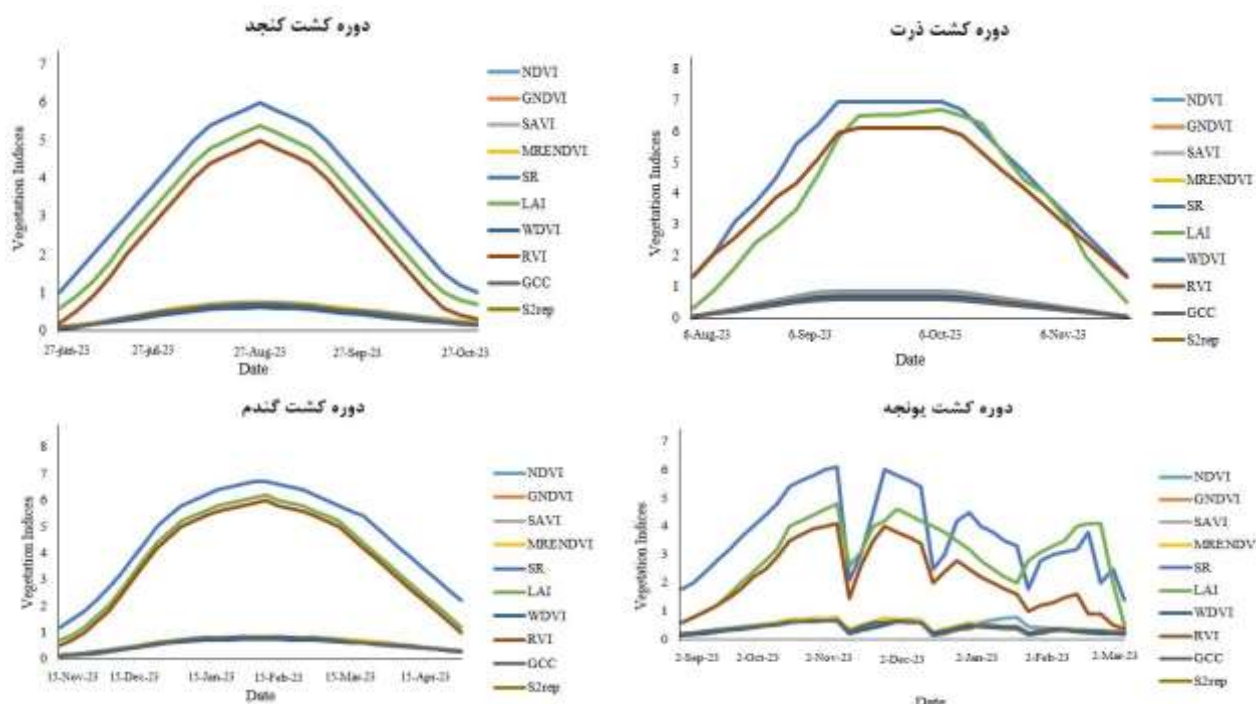
برای داده‌های سنتینل-۲، باندهای طیفی واقع در محدوده‌های مرئی، مادون قرمز نزدیک و مادون قرمز میانی شامل باندهای ۲ تا ۷، ۸A، ۸B، ۱۱ و ۱۲ با قدرت تفکیک مکانی ۱۰، ۲۰ و ۶۰ متر انتخاب شدند. به منظور یکسان‌سازی تفکیک مکانی باندها و فراهم‌سازی امکان ترکیب آن‌ها با داده‌های راداری سنتینل-۱، کلیه باندها به تفکیک مکانی ۱۰ متر بازنمونه‌گیری^۱ شدند. در این فرآیند از روش درون‌یابی دوخطی^۲ استفاده گردید، زیرا این روش به‌طور ویژه برای داده‌های پیوسته نظیر بازتاب طیفی مناسب بوده و منجر به نمایش نرم‌تر تغییرات مکانی و کاهش ناهمگنی‌های پیکسلی می‌شود (Muhammad Zayyanul & Projo, 2019; Trinder & Salah, 2012).

تصاویر سنتینل-۲ پس از طی مراحل تصحیحات پیش‌پردازشی شامل تصحیح جوئی، تبدیل درخشندگی به بازتاب سطح زمین و اعمال ماسک ابر پردازش شدند. برای ماسک‌گذاری ابر از الگوریتم Fmask استفاده گردید و مقادیر پیکسل‌های ابری با بهره‌گیری از درون‌یابی زمانی میان تصاویر پیش و پس از رخداد ابر اصلاح شدند (Oishi et al., 2018). در ادامه، مجموعه‌ای از شاخص‌های طیفی گیاهی محاسبه و به‌صورت لایه‌های مجزا بر روی ترکیب تصاویر مربوط به دوره‌های کشت پاییزه و تابستانه اعمال شد. این شاخص‌ها نقش مهمی در پایش فنولوژی گیاهان و افزایش دقت طبقه‌بندی محصولات کشاورزی ایفا کردند. همچنین، الگوهای فنولوژیکی محصولات مورد مطالعه در شکل ۲ نمایش داده شده‌اند.

1 Google Earth Engine (GEE)

2 resampling

3 bilinear interpolation



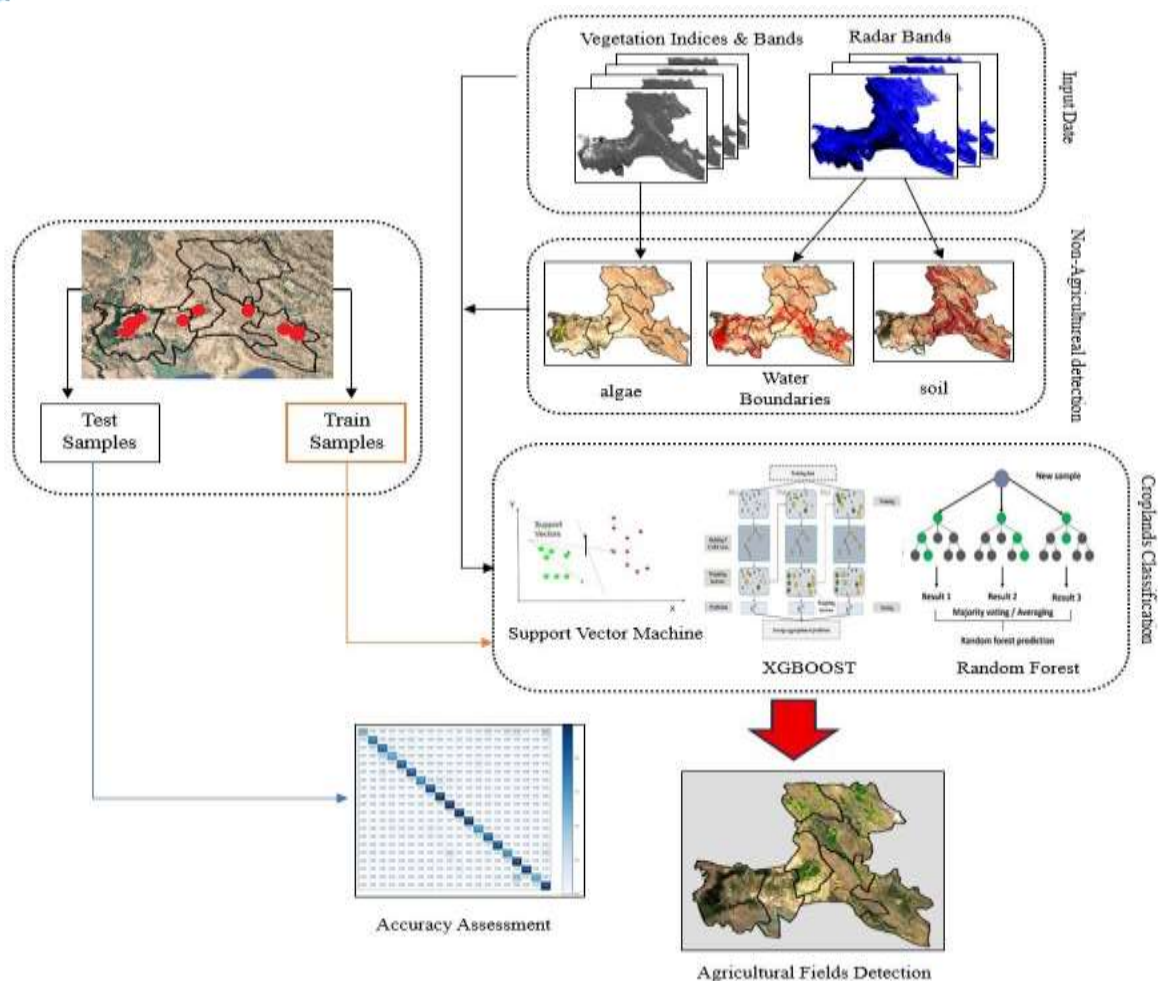
شکل ۲. تغییرات فنولوژی محصولات مورد مطالعه ۲۰۲۳

برای داده‌های سنتینل ۱، تصاویر مربوط به مد (Interferometric Wide (IW) و مسیر Descending فراخوانی شدند و از دو پلاریزاسیون VV و VH به منظور استخراج ویژگی‌های راداری استفاده گردید. داده‌های سنتینل ۱ در محیط گوگل ارث انجین به صورت پیش فرض شامل تصحیحات هندسی و رادیومتریکی بوده و ضرایب پس‌پراکنش (σ^0) آن‌ها در مقیاس دسی‌بل ارائه می‌شوند. به منظور بهبود کیفیت داده‌ها و کاهش نویز اسپکل ذاتی تصاویر راداری، فیلتر تطبیقی لی‌بر روی باندهای VV و VH اعمال شد. در این فیلتر از پنجره‌ای با ابعاد 5×5 پیکسل استفاده گردید که با بهره‌گیری از آمار محلی (میانگین و واریانس) وزن‌دهی تطبیقی انجام می‌دهد. این فرآیند ضمن حفظ جزئیات و مرزهای عوارض سطحی، موجب حذف داده‌های پرت و افزایش پایداری رادیومتریکی تصویر گردید.

سپس، نسبت VV/VH محاسبه شد تا با داده‌های اپتیکی سنتینل ۲ و شاخص‌های طیفی مربوطه تلفیق گردد. این اقدام، با هدف افزایش توان تفکیک بین انواع پوشش‌های گیاهی و اراضی کشاورزی، دقت طبقه‌بندی نهایی را به طور قابل توجهی بهبود بخشید (Hasan et al., 2021; Lavreniuk et al., 2017; Rubel et al., 2021).

در مرحله بعد، داده‌های ترکیبی حاصل از باندهای اپتیکی، شاخص‌های گیاهی (جدول ۴) و ویژگی‌های راداری به عنوان ورودی مدل‌های یادگیری ماشین مورد استفاده قرار گرفتند. نمونه‌های آموزشی استخراج‌شده برای محصولات مختلف کشاورزی، به دو مجموعه آموزشی و آزمون تقسیم شدند و طبقه‌بندی نهایی با بهره‌گیری از مدل‌های جنگل تصادفی (RF)، ماشین بردار پشتیبان (SVM) و گرادیان تقویتی شدید (XGB) انجام شد. برای آموزش مدل طبقه‌بندی، از داده‌های نمونه میدانی مربوط به کلاس‌های مختلف کاربری/پوشش زمین شامل گندم، یونجه، کجند، ذرت، خاک بدون پوشش و پهنه‌های آبی و جلبک و خزہ سطح آب استفاده شد. این نمونه‌ها در قالب فایل‌های CSV حاوی مختصات جغرافیایی و برچسب کلاس آماده‌سازی گردید و به کمک کتابخانه Pandas و توابع پردازشی، به FeatureCollection در گوگل ارث انجین، تبدیل شدند. داده‌های نمونه به نحوی انتخاب شده‌اند که تنوع مکانی و پوشش گیاهی را در سطح منطقه به خوبی پوشش دهند. سپس داده‌های آموزشی با نمونه‌برداری از تصویر ترکیبی به مقیاس مکانی ۱۰ متر استخراج و برای آموزش مدل‌های یادگیری ماشین مورد استفاده قرار گرفتند و در نهایت مدل‌ها مورد ارزیابی قرار گرفتند.

در نهایت، نقشه‌های طبقه‌بندی تولید گردید و برای محصولات مورد بررسی، کلاس‌های مربوطه استخراج و مساحت کشت در سطح کل منطقه و به تفکیک هر شهر محاسبه شد. این نتایج علاوه بر نمایش الگوی مکانی کشت محصولات، امکان تحلیل دقیق‌تر وضعیت کشاورزی منطقه را فراهم ساختند. روند انجام کار در شکل ۳، نشان داده شده است.



شکل ۳. فلوچارت مراحل انجام کار

داده‌های زمینی

در سال‌های زراعی ۱۴۰۳-۱۴۰۰، داده‌های میدانی از مزارع کشاورزی در شهرستان‌های زیرحوضه مارون-جراحی خوزستان جمع‌آوری شد. موقعیت مزارع پایش شده به صورت میدانی توسط کارشناسان در شکل ۲ با دایره‌های قرمز مشخص شده است. بر اساس داده‌های گردآوری‌شده، محصولات اصلی منطقه شامل گندم، ذرت، یونجه، کنجد هستند. در جدول ۲، مزارع منتخب حاصل از بازدیدهای میدانی، با بالاترین مساحت زیرکشت نشان داده شده است.

جدول ۲. مشخصات کلی پایلوت‌های منتخب

شهرستان	نوع محصول	Y	X
بهبهان	گندم	۳۳۸۹۰۹۸	۴۳۱۹۷۰
امیدیه		۳۴۲۰۶۹۵	۳۸۸۵۶۴
رامشیر		۳۴۱۷۷۶۹	۳۳۸۸۷۸
شادگان		۳۳۹۶۶۳۹	۲۷۸۸۱۰
بندرماهشهر		۳۴۰۶۷۹۸	۳۳۳۴۸۳
بهبهان	ذرت	۳۳۹۴۸۱۰	۴۲۶۱۱۶
امیدیه		۳۴۲۰۲۸۳	۳۸۹۲۱۸
رامشیر	کنجد	۳۴۱۶۵۸۸	۳۴۰۵۱۱
بهبهان	یونجه	۳۳۹۵۶۷۲	۴۲۵۴۹۶
رامشیر		۳۴۱۶۵۸۵	۳۴۰۳۰۵

برای انجام طبقه‌بندی نظارت‌شده محصولات کشاورزی، نمونه‌های آموزشی و آزمون از داده‌های میدانی انتخاب شدند. نمونه‌های آموزشی صرفاً برای ساخت مدل اولیه طبقه‌بندی مورد استفاده قرار گرفتند، در حالی که نمونه‌های آزمون برای ارزیابی دقت، صحت و قابلیت اعتماد روش‌های طبقه‌بندی به کار رفتند. در این مطالعه، ۷۰٪ از داده‌های میدانی به صورت تصادفی انتخاب و به عنوان داده‌های آموزشی مورد استفاده قرار گرفت و بخش باقی‌مانده برای آزمون اختصاص یافت. جزئیات مربوط به تعداد نمونه‌های آموزشی هر کلاس محصول در جدول ۳، ارائه شده است.

جدول ۳. تعداد نقاط آموزشی برای هر محصول در منطقه مطالعاتی

ذرت	کنجد	یونجه	گندم	کلاس گیاهی
۶۵۰	۵۰۰	۶۲۰	۶۳۵	کل نمونه ها
۴۵۵	۳۵۰	۴۳۴	۴۴۴	نمونه‌های آموزشی
۰/۲۷	۰/۲۱	۰/۲۶	۰/۲۵	IR= ۱/۳

به منظور بررسی تعادل میان تعداد نمونه‌های هر کلاس، از معیار درجه عدم توازن استفاده شد (Kumar et al., 2022). در این روش، ابتدا سهم هر کلاس از کل نمونه‌ها محاسبه شد. سپس شاخص نسبت عدم توازن طبق رابطه ۱ زیر محاسبه گردید:

$$IR = \frac{\max \{ni\}}{\min \{ni\}} \quad (\text{رابطه ۱})$$

ni: تعداد نمونه‌های کلاس i

محاسبه شاخص‌های طیفی و راداری

به منظور تفکیک بین کلاس‌های مختلف زراعی و باغی و سطوح غیر زراعی علاوه بر استفاده از شاخص‌های متداول پوشش گیاهی، مجموعه‌ای از شاخص‌های ترکیبی مبتنی بر باندهای قرمز لبه و نسبت‌های طیفی خاص گیاهان گرمسیری طراحی و به کار گرفته شد. همچنین تلفیق همزمان داده‌های راداری و شاخص‌های اپتیکی به بهبود شناسایی محصولات با پوشش متراکم و تفکیک بهتر کلاس‌های دارای رفتار طیفی مشابه می‌گردد. در جدول ۴، ویژگی‌های هر شاخص نشان داده شده است.

جدول ۴. معرفی شاخص‌های طیفی

شاخص طیفی	رابطه محاسباتی	کاربرد	منبع
NDVI	(NIR-RED)/(NIR+RED)	تراکم و سرسبزی گیاهان را اندازه‌گیری می‌کند	Rokni & Musa, (2019)
GNDVI	(NIR-Green)/(NIR+Green)	نسخه‌ای از NDVI است که به جای نور قرمز از نور سبز استفاده می‌کند تا حساسیت آن به کلروفیل گیاهی افزایش یابد	(Nițu et al., 2025)
SAVI	$((NIR - RED) / (NIR + RED + 0.5)) * 1.5$	در نواحی با پوشش گیاهی پراکنده کاربرد دارد و می‌تواند تصحیح خاک را بهبود بخشد	Rondeaux et al., (1996)
MRENDVI	(Rededge3-Rededge1)/(Rededge3+Rededge1)	این شاخص بر حساسیت به تغییرات جزئی در تراکم برگ‌ها و مراحل رشد گیاه تمرکز دارد و در پایش هوشمند کشاورزی و تشخیص استرس‌های گیاهی کاربرد دارد	(Gamal et al., 2020)
S2rep	$((2-Red\ Edge\ 1/(Red+Red\ Edge\ 3))/(Red\ Edge\ 2-Red\ Edge\ 1))*35+705$	تغییرات S2REP با غلظت کلروفیل رابطه مستقیم دارد. با افزایش کلروفیل، موقعیت لبه قرمز به طول موج‌های بلندتر منتقل می‌شود.	(Ali et al., 2022)
SR	$NIR / (RED + 1e-10)$	بررسی تغییرات کلروفیل و وضعیت نیتروژن در گیاه	(Yang et al., 2022)
LAI	$3.618 * EVI - 0.118$	کمیتی فیزیکی است که به ازای هر واحد سطح زمین، مقدار مساحت یک‌طرفه برگ‌ها را اندازه‌گیری می‌کند.	(Nițu et al., 2025)
RVI	NIR/RED	این شاخص به عنوان اندازه‌گیری غیرخطی سرسبزی عمل می‌کند	Baugh & (Groeneveld, 2006)
WDVI	$NIR-Slope*Red\ Slope = nir\ soil / red\ soil$	برای کاهش تأثیر روشنایی خاک در اندازه‌گیری پوشش گیاهی توسعه یافته است	(Clevers, 1989)
GCC	Green/(Red+Green+Blue)	تغییرات سبزی برگ‌ها را در طول فصل رشد ثبت می‌کند و برای تشخیص سبز شدن، اوج رشد و رنگ‌پریدگی پاییزه کاربرد دارد.	(Liu et al., 2022)

لازم به ذکر است که در مرحله پیش‌پردازش داده‌ها، به جای استفاده از تنها یک مجموعه شاخص عمومی، مجموعه‌ای از شاخص‌های



اختصاصی برای ذرت شامل $MCARI$ ، $OSAVI$ و RVI طراحی و به مدل افزوده شد که تأثیر چشمگیری در بهبود دقت شناسایی و تفکیک این محصول از سایر کلاس‌ها داشت. به علاوه، تلفیق سه متغیر راداری شامل VV ، VH و VV/VH نسبت به شاخص‌های اپتیکی، موجب افزایش حساسیت مدل‌ها نسبت به رطوبت خاک و ساختار گیاهی گردید.

مدل‌های یادگیری ماشین

جنگل تصادفی

مدل جنگل تصادفی (RF) از روش‌های قدرتمند و کاربردی در دسته‌ی مدل‌های یادگیری نظارت‌شده به‌شمار می‌رود (Breiman, 2001; Goldblatt et al., 2016). این مدل نخستین بار توسط (Breiman, 2001) معرفی شد. این مدل به‌عنوان مجموعه‌ای از درخت‌های تصمیم تعریف می‌شود که در آن هر درخت به طبقه‌ی یک نمونه رأی می‌دهد و طبقه‌ای که بیشترین رأی را کسب کند به‌عنوان خروجی نهایی انتخاب می‌شود (Sun & Schulz, 2015). این مدل توانایی بالایی در پردازش داده‌های با ابعاد زیاد داشته و نسبت به بیش‌برازش نیز مقاوم است (Breiman, 2001). همچنین، استفاده از تصادفی‌سازی در انتخاب نمونه‌ها و ویژگی‌ها موجب می‌شود مدل وابستگی کمتری به داده‌های خاص داشته باشد و در نتیجه عملکرد بهتری روی داده‌های جدید نشان دهد.

ترکیب خروجی درخت‌های تصمیم متعدد در این مدل، اثر نویز موجود در داده‌ها را کاهش داده و پیش‌بینی‌هایی پایدارتر و قابل اعتمادتر ارائه می‌کند. علاوه بر این، جنگل تصادفی توانایی پردازش داده‌های پیچیده و غیرخطی را داشته و می‌تواند با مجموعه‌های داده‌ای شامل تعداد زیادی ویژگی یا حتی داده‌های ناقص به‌خوبی کار کند.

از مزایای این روش می‌توان به سهولت پیاده‌سازی و کاربردهای گسترده‌ی آن در حوزه‌های مختلف مانند طبقه‌بندی و رگرسیون اشاره کرد (Ray, 2019). این مدل در زمینه‌هایی همچون شناسایی الگوها (مانند تحلیل تصاویر)، پیش‌بینی مقادیر عددی (مانند برآورد عملکرد محصولات کشاورزی) و مسائل متنوع علمی به‌طور گسترده مورد استفاده قرار می‌گیرد.

ترکیب دقت بالا، تعمیم‌پذیری مناسب، و مقاومت در برابر نویز و داده‌های ناقص سبب شده است که جنگل تصادفی به یکی از پرکاربردترین و قابل اعتمادترین مدل‌های یادگیری ماشین تبدیل شود (رضائی اعتدالی، ه. و احمدی، م. ۱۴۰۳).

در این مطالعه، مدل RF با استفاده از Google Earth Engine Python API اجرا شد. علاوه بر پیاده‌سازی الگوریتم استاندارد جنگل تصادفی، تغییرات مشخصی در ساختار پارامترهای آن اعمال گردید تا عملکرد مدل برای طبقه‌بندی محصولات زراعی در منطقه گرم و نیمه‌خشک مارون بهینه شود. در این راستا، پارامترهای کلیدی شامل تعداد درخت‌ها ($numberOfTrees = 100$)، تعداد ویژگی‌های انتخاب شده در هر گره ($variablesPerSplit = 5$)، حداقل تعداد نمونه در برگ‌ها ($minLeafPopulation = 3$)، نسبت نمونه‌های انتخابی در هر درخت ($bagFraction = 0.7$) و مقدار $Seed = 42$ برای تولید نتایج قابل تکرار بودند. انتخاب این پارامترها بر اساس آزمایش‌های اولیه روی داده‌های آموزشی انجام شد تا تعادل بین دقت و جلوگیری از بیش‌برازش حاصل شود. به‌طور خاص، افزایش تعداد درخت‌ها باعث کاهش واریانس مدل و پایداری پیش‌بینی‌ها شد، اما بیش از ۱۰۰ درخت بهبود قابل توجهی ایجاد نکرد. مقدار ۵ برای ویژگی‌های انتخاب شده در هر گره تنوع کافی بین درخت‌ها ایجاد کرده و از پیچیدگی بیش‌ازحد جلوگیری کرد. حداقل نمونه در برگ‌ها مانع از ایجاد برگ‌های بسیار کوچک و بیش‌برازش شد و نسبت نمونه‌های انتخابی ($bagFraction = 0.7$) باعث شد هر درخت روی زیرمجموعه متفاوتی از داده‌ها آموزش ببیند و وابستگی به نمونه‌های خاص کاهش یابد.

ماشین بردار پشتیبان

مدل ماشین بردار پشتیبان (SVM) در مطالعات متنوعی از جمله در حوزه سنجش از دور مورد بررسی قرار گرفته است (George et al., 2006; Marcinkowska et al., 2014; Pal & Mather, 2014). این مدل نخستین بار توسط (Vapnik, 2013) معرفی شد. مدل SVM به دنبال یافتن ابرصفحه‌ی بهینه در فضای n -بعدی است که بیشترین حاشیه جداسازی را میان طبقات ایجاد کند.

این روش با استفاده از توابع کرنل، از جمله کرنل پایه شعاعی (RBF)، قادر است مرزهای تصمیم‌گیری غیرخطی را نیز مدیریت کند. از جمله پارامترهای کلیدی این مدل می‌توان به C (پارامتر جریمه)، γ (ضریب کرنل) و نوع کرنل اشاره کرد (Khosravi et al., 2021). گاما نقشی کلیدی در کنترل انعطاف‌پذیری مرز تصمیم‌گیری دارد؛ به‌طوری که مقادیر پایین‌تر گاما به مرزهای ساده‌تر و تعمیم‌پذیرتر منجر می‌شوند، درحالی که مقادیر بالاتر آن می‌تواند باعث ایجاد مرزهای بسیار پیچیده و احتمال بیش‌برازش گردد. در این تحقیق، با توجه به دقت بالاتر کرنل RBF در مقایسه با سایر کرنل‌ها، از این نوع کرنل استفاده شده و مقدار گاما برابر با 0.0002 در نظر گرفته شد (Alejo

(et al., 2006). در این تحقیق، با توجه به دقت بالاتر کرنل RBF در مقایسه با سایر کرنل‌ها، از این نوع کرنل استفاده شده و مقدار گاما برابر 0.0002 تا 0.0004 و مقدار C در بازه ۱ تا ۱۰۰ در نظر گرفته شد.

بیشینه‌سازی حاشیه در ماشین بردار پشتیبان سبب ایجاد تغییراتی در تعیین مرز تصمیم‌گیری می‌شود و این ویژگی امکان پردازش داده‌های پیچیده و چندبعدی را فراهم می‌آورد. به همین دلیل، ماشین بردار پشتیبان یکی از روش‌های کارآمد برای تحلیل داده‌های متنوع و غیرخطی به‌شمار می‌آید (Ghodsi et al., 2021a).

XGBoost

به عنوان یکی از مدل‌های دقیق یادگیری ماشین شناخته می‌شود که به‌طور ویژه در تحلیل داده‌های حجیم کارایی بالایی دارد. این مدل بر پایه‌ی روش تقویت گرادیان شدید توسعه یافته و با به‌کارگیری مجموعه‌ای از مدل‌های ساده مانند درخت‌های تصمیم، قادر است پیش‌بینی‌هایی دقیق و قابل اعتماد ارائه دهد (Sahin, 2020). همچنین، در پردازش داده‌های پراکنده توانمند بوده و در مواجهه با متغیرهای ورودی متنوع و پیچیده، ابزاری کاربردی محسوب می‌شود (Chen & Guestrin, 2016). XGBoost در چارچوب تقویت گرادیانی (Gradient Boosting) عمل می‌کند، به‌طوری‌که در هر مرحله، یک درخت تصمیم جدید براساس خطاهای باقی‌مانده‌ی درخت قبلی ساخته می‌شود. این فرایند تکرارشونده به مدل امکان می‌دهد پیش‌بینی‌های خود را اصلاح کرده و در نهایت خروجی با دقت بالاتری دست یابد.

اجرای مدل XGBoost در محیط Python با بسته xgboost انجام شد. پارامترهای کلیدی شامل تعداد درخت‌ها (n_estimators)، عمق حداکثر درخت‌ها (max_depth)، نرخ یادگیری (learning_rate) و پارامترهای تنظیم‌کننده مانند gamma، subsample و colsample_bytree بودند. بهینه‌سازی این پارامترها با جستجوی شبکه‌ای (Grid Search) و اعتبارسنجی متقاطع پنج‌تایی (5-fold Cross-Validation) انجام شد تا ترکیبی از پارامترها انتخاب گردد که کمترین خطا و بیشترین تعمیم‌پذیری روی داده‌های آزمون را ارائه دهد. انتخاب این روش بهینه باعث شد که مدل هم دقت بالایی داشته باشد و هم از بیش‌برازش جلوگیری شود، و در نتیجه نتایج قابل اعتماد و پایدار ارائه کند.

توابع کتابخانه‌ای مورد استفاده

تمامی پردازش‌ها و مدل‌سازی در محیط Python و با استفاده از Google Earth Engine Python API انجام شد. برای این منظور از کتابخانه‌های زیر استفاده گردید:

geemap: به‌عنوان رابط برنامه‌نویسی و تجسم داده‌های ماهواره‌ای در Google Earth Engine، امکان پردازش تصاویر ستینل، محاسبه شاخص‌های طیفی، نمایش نقشه‌ها و افزودن لایه‌ها و برچسب‌ها را فراهم می‌کند.

(Earth Engine Python API): برای دسترسی به داده‌های ماهواره‌ای با وضوح بالا و انجام پردازش‌های حجیم روی سرورهای GEE به‌کار گرفته شد. این کتابخانه برای فیلترکردن تصاویر بر اساس تاریخ و موقعیت مکانی، محاسبه شاخص‌های طیفی، نمونه‌برداری از داده‌های آموزشی، آموزش و ارزیابی مدل‌های یادگیری ماشین و محاسبه مساحت محصولات کشاورزی استفاده شد.

pandas: برای مدیریت و پردازش داده‌های جدولی و CSV شامل نمونه‌های زمینی، اصلاح نام ستون‌ها و استخراج مقادیر طول و عرض جغرافیایی استفاده شد.

ast: جهت تبدیل رشته‌های JSON-like مختصات جغرافیایی موجود در داده‌های CSV به ساختار دیکشنری Python برای استخراج طول و عرض جغرافیایی مورد استفاده قرار گرفت.

geopandas: برای پردازش داده‌های مکانی و Shapefile، شامل خواندن محدوده شهرستان‌ها و رودخانه‌ها، تغییر سیستم مختصات، محاسبه مرکز هندسی (centroid) برای برچسب‌گذاری و تبدیل FeatureCollection به GeoDataFrame. json: برای تبدیل داده‌ها به فرمت GeoJSON جهت تعامل با Earth Engine و geemap و انتقال داده‌ها بین Python و GEE

ارزیابی مدل‌ها

به‌منظور ارزیابی عملکرد مدل‌های طبقه‌بندی به‌کاررفته در این پژوهش، ابتدا ماتریس ابهام^۱ تهیه شد و سپس شاخص‌های زیر محاسبه



گردید: صحت کلی، ضریب کاپا، صحت تولیدکننده و صحت کاربر^۴

ماتریس ابهام

ماتریس ابهام یا «ماتریس خطا» جدولی است که مقدار نمونه‌های پیش‌بینی‌شده برای هر کلاس را در مقابل مقدار واقعی آن‌ها نشان می‌دهد. در این ماتریس، سطرها معمولاً کلاس‌های پیش‌بینی‌شده و ستون‌ها کلاس‌های واقعی را نمایندگی می‌کنند و سلول‌های قطر اصلی (i,i) نشان‌دهنده تعداد نمونه‌هایی هستند که به درستی طبقه‌بندی شده‌اند (Nicolau et al., 2024). و با شناسایی و تجزیه و تحلیل الگوی اشتباهات طبقه‌بندی مشخص می‌سازد کدام کلاس‌ها با هم مخلوط شده‌اند (Feizizadeh et al., 2022a). رابطه ۲، ماتریس درهم‌ریختگی را نشان می‌دهد.

Confusion Matrix=

$$\begin{bmatrix} \text{False Positive (FP)} & \text{True Positive (TP)} \\ \text{True Negative (TN)} & \text{False Negative (FN)} \end{bmatrix}$$

(رابطه ۲)

یک ماتریس درهم‌ریختگی دارای چهار جزء اصلی است که عبارتند از: True Positive (TP): تعداد نمونه‌هایی که به درستی در کلاس مربوطه پیش‌بینی شده‌اند، False Positive (FP): تعداد نمونه‌هایی که به نادرست در کلاس دیگری پیش‌بینی شده‌اند، False Negative (FN): تعداد نمونه‌هایی که به نادرست در کلاس مربوطه پیش‌بینی نشده‌اند. True Negative (TN): تعداد نمونه‌هایی که به درستی در کلاس نادرست پیش‌بینی شده‌اند.

صحت کلی

صحت کلی یک معیار ارزیابی است که برای بررسی صحت نتایج طبقه‌بندی به کار می‌رود. این معیار برابر نسبت تعداد نمونه‌های به درستی طبقه‌بندی‌شده به کل نمونه‌های بررسی‌شده است. به طور عملی، صحت کلی را می‌توان با جمع مقادیر قطر اصلی ماتریس درهم‌ریختگی و تقسیم بر مجموع کل نمونه‌ها محاسبه کرد (Aziz et al., 2024). صحت کلی به صورت رابطه ۳، تعریف می‌شود.

$$\text{Overall Accuracy} = \frac{N_{\text{Correct}}}{N_{\text{Total}}}$$

(رابطه ۳)

N_{Correct} : تعداد نمونه‌های که به درستی طبقه‌بندی شده‌اند، N_{Total} : کل تعداد نمونه‌ها است.

Precision

معیاری برای ارزیابی عملکرد مدل‌های طبقه‌بندی می‌باشد. این شاخص نشان می‌دهد از میان نمونه‌هایی که مدل آن‌ها را به عنوان مثبت (یا متعلق به یک کلاس مشخص) پیش‌بینی کرده است، چند درصد واقعاً متعلق به آن کلاس هستند (Foody, 2023).

$$\text{Precision} = \frac{TP}{TP+FP}$$

(رابطه ۴)

True Positives: تعداد نمونه‌هایی که واقعاً به کلاس مثبت تعلق داشتند و مدل نیز آن‌ها را مثبت پیش‌بینی کرده بود.

False Positives: تعداد نمونه‌هایی که مدل آن‌ها را به کلاس مثبت پیش‌بینی کرده بود، اما واقعاً به آن کلاس تعلق نداشته‌اند.

Recall

یکی از معیارهای مهم ارزیابی مدل‌های طبقه‌بندی است. این شاخص نشان می‌دهد از میان تمام نمونه‌هایی که واقعاً به کلاس هدف تعلق دارند، مدل چه نسبت از آن‌ها را به درستی تشخیص داده است.

$$\text{Recall} = \frac{TP}{TP+FN}$$

(رابطه ۵)

FN: تعداد نمونه‌هایی است که واقعاً به کلاس هدف تعلق داشتند اما مدل آن‌ها را از دست داده است (Foody, 2023).

F1-Score

به منظور سنجش عملکرد مدل‌های طبقه‌بندی، علاوه بر شاخص‌های رایج نظیر دقت (Precision) و بازیابی (Recall) از معیار F1-Score بهره گرفتیم. دقت معرف نسبت نمونه‌هایی است که مدل مثبت پیش‌بینی کرده و واقعاً مثبت بودند، و بازیابی بیان‌کننده نسبت نمونه‌های واقعاً مثبت است که توسط مدل به درستی شناسایی شده‌اند (Zheng et al., 2024).

1 Overall Accuracy

2 Kappa Coefficient

3 Producer's Accuracy

4 User's Accuracy

معیار F1-Score به صورت میانگین هارمونیک دقت و بازیابی تعریف می‌شود:

$$F1 = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (\text{رابطه ۶})$$

ضریب کاپا^۱

ضریب کاپا معیاری آماری برای سنجش میزان توافق بین نتایج طبقه‌بندی و داده‌های مرجع است که تاثیر توافق تصادفی را کاهش می‌دهد (Feizizadeh et al., 2022b). مقدار کاپا بین -۱ و +۱ نوسان می‌کند، به گونه‌ای که $\kappa=1$ نشان‌دهنده توافق کامل، $\kappa=0$ نشان‌دهنده توافق معادل با شانس و مقادیر منفی نشان‌دهنده توافق کمتر از شانس است.

$$\text{Kappa} = \frac{P_0 - P_e}{1 - P_e} \quad (\text{رابطه ۷})$$

طبق رابطه ۷، در آن: P_0 برابر است با تعداد نمونه‌های به درستی طبقه‌بندی شده تقسیم بر کل تعداد نمونه‌ها و P_e برابر است با مجموع حاصل ضرب‌های تعداد واقعی نمونه‌ها و تعداد پیش‌بینی شده در هر دسته، تقسیم بر مجذور کل تعداد نمونه‌ها (Zhan et al., 2024). به منظور اعتبارسنجی فضایی داده‌ها، در این پژوهش از رویکرد اعتبارسنجی تصادفی لایه‌بندی شده در سطح مکانی استفاده شد؛ به این معنا که داده‌های نمونه از چند شهرستان مستقل در محدوده حوضه مارون-جراحی برداشت شدند و در فرآیند اعتبارسنجی، تفکیک آموزش و آزمون به گونه‌ای انجام شد که هر شهرستان به صورت مستقل در مجموعه داده لحاظ گردد (randomColumn با seed ثابت و تقسیم ۷۰/۳۰). به این ترتیب، هرچند از روش مرسوم spatial cross-validation به صورت رسمی استفاده نشد، اما طراحی نمونه‌ها و جداسازی مکانی شهرستان‌ها در داده‌های آموزشی و آزمایشی موجب شد ارزیابی مدل نسبت به تغییرات مکانی حساس باشد و به شکل مؤثری اثر وابستگی فضایی کاهش یابد.

یافته‌های پژوهش و بحث

بررسی عملکرد SVM

بر اساس شکل ۴، تحلیل ماتریس ابهام نشان داد که مدل در تفکیک کلاس‌های زراعی عملکرد نسبتاً متوسطی داشته است، چراکه میزان دقت در میان کلاس‌ها یکسان نبوده و تفاوت‌هایی در نحوه شناسایی هر محصول مشاهده می‌شود. در مجموع، بیشترین تعداد نمونه‌های درست طبقه‌بندی شده مربوط به کلاس‌های گندم و یونجه است، در حالی که برخی کلاس‌ها مانند کنجد و ذرت دچار تداخل و اشتباه در شناسایی شده‌اند. این الگو بیانگر آن است که در کلاس‌هایی که ویژگی‌های طیفی و ساختاری گیاهان در طول فصل رشد به خوبی از یکدیگر متمایز می‌شوند، مدل قادر به تشخیص صحیح‌تر بوده است؛ اما در مواردی که رفتار طیفی مشابه یا هم‌پوشانی فنولوژیکی وجود داشته، افت دقت مشاهده می‌شود.

به طور خاص، مشاهده می‌شود که در کلاس ذرت تعدادی از نمونه‌ها با خاک و جلبک اشتباه گرفته شده‌اند، که می‌تواند ناشی از بازتاب‌های طیفی مشابه در باندهای خاص و یا تأثیر رطوبت خاک و پوشش گیاهی کم تراکم باشد. در مورد کنجد، اشتباه در طبقه‌بندی با کلاس‌های یونجه و گندم به طور قابل توجهی دیده می‌شود. این امر ممکن است به دلیل نزدیکی مراحل رشد، تراکم پوشش گیاهی و رنگ مشابه تاج پوشش در بازه‌های زمانی برداشت داده‌ها باشد. در مقابل، کلاس خاک با وجود تنوع طیفی بالا، دقت مناسبی در تفکیک از سایر کلاس‌ها نشان داده است که بیانگر توانایی داده‌های مورد استفاده در بازنمایی ویژگی‌های طیفی سطح خاک است.

نکته قابل توجه دیگر مربوط به طبقه مرز آبی و جلبک است که بخشی از نمونه‌های آن‌ها به صورت متقابل با یکدیگر اشتباه گرفته شده‌اند. چنین تداخلی احتمالاً به پدیده اختلاط پیکسل‌ها در نواحی حاشیه‌ای و یا تغییرات ناگهانی در بازتاب سطح آب و حضور مواد آلی در سطح آن مربوط است. این موضوع نشان می‌دهد که در مرز بین کلاس‌های آبی و خشکی، نیاز به دقت مکانی بالاتر و یا به کارگیری شاخص‌های طیفی ویژه برای تفکیک این نواحی وجود دارد.

بر اساس جدول ۵، نتایج شاخص‌های ارزیابی دقت نشان می‌دهد، در برخی کلاس‌ها اختلاف میان دقت و بازخوانی قابل توجه بوده

1 Kappa coefficient

2 precision

3 Recall

است. اما، به طور کلی، مقادیر به دست آمده در بازه‌ای میان ۰/۴۵ تا ۰/۷۱ قرار دارند که بیانگر دقت متوسط تا نسبتاً مطلوب در سطح کلی است. با این حال، در برخی کلاس‌ها اختلاف میان دقت و بازخوانی قابل توجه بوده و این امر می‌تواند نشان‌دهنده جهت‌دار بودن خطاهای مدل در تشخیص نمونه‌های آن کلاس‌ها باشد. در میان کلاس‌ها، مرز آبی بالاترین مقدار دقت (۰/۷۱) را داشته است که نشان می‌دهد مدل توانسته است در شناسایی صحیح این نواحی عملکرد بهتری داشته باشد. این امر را می‌توان به ویژگی‌های طیفی متمایز سطوح آبی نسبت داد که معمولاً در محدوده‌های خاصی از بازتاب طیفی قرار می‌گیرند و از سایر کلاس‌ها قابل تفکیک‌اند. با این حال، مقدار بازخوانی پایین‌تر این کلاس (۰/۵۴) نشان می‌دهد که بخشی از نمونه‌های واقعی مرز آبی به درستی شناسایی نشده و در طبقات دیگر، به ویژه جلبک یا خاک، قرار گرفته‌اند. این پدیده احتمالاً ناشی از وجود پیکسل‌های مختلط در نواحی حاشیه‌ای آب و خشکی یا تغییرات جزئی در عمق و شفافیت آب است.

در مقابل، کلاس کنگد پایین‌ترین مقدار دقت (۰/۴۵) را داشته است، که حاکی از آن است که بخشی از پیکسل‌هایی که به عنوان کنگد شناسایی شده‌اند، در واقع به کلاس‌های دیگر تعلق داشته‌اند. با این حال، بازخوانی نسبتاً بالاتر (۰/۵۴) در همین کلاس نشان می‌دهد که درصد قابل توجهی از نقاط واقعی کنگد نیز به درستی شناسایی شده‌اند. چنین الگویی معمولاً زمانی رخ می‌دهد که ویژگی‌های طیفی محصول با برخی کلاس‌های دیگر (مانند ذرت یا یونجه) هم‌پوشانی دارد و باعث می‌شود مدل در هنگام تشخیص، محدوده‌ی طیفی مشابه را به چند کلاس مختلف نسبت دهد.

کلاس گندم با دقت ۰/۶۵ و بازخوانی ۰/۵۴ عملکرد متعادلی نشان داده است. این نتیجه بیانگر آن است که مدل در شناسایی صحیح گندم نسبت به سایر محصولات عملکرد بهتری دارد، که احتمالاً ناشی از ساختار پوشش گیاهی منظم، تراکم بالای گیاه و بازتاب‌های طیفی مشخص در باندهای مرئی و فروسرخ نزدیک است. با این وجود، تفاوت میان مقدار دقت و بازخوانی در این کلاس می‌تواند به خطاهای جزئی در تمایز گندم از محصولات دارای مراحل رشدی مشابه، به ویژه یونجه، مربوط باشد.

کلاس‌های یونجه و جلبک نیز الگوی مشابهی از نظر عملکرد نشان داده‌اند، به گونه‌ای که مقادیر دقت و بازخوانی هر دو در محدوده ۰/۵ تا ۰/۵۵ قرار دارند. این میزان از دقت متوسط را می‌توان ناشی از تنوع در ساختار و شرایط زیستی این پوشش‌ها دانست. برای مثال، در مناطق با رشد غیریکنواخت یا پوشش پراکنده، بازتاب‌های طیفی می‌تواند در محدوده کلاس‌های مجاور قرار گیرد.



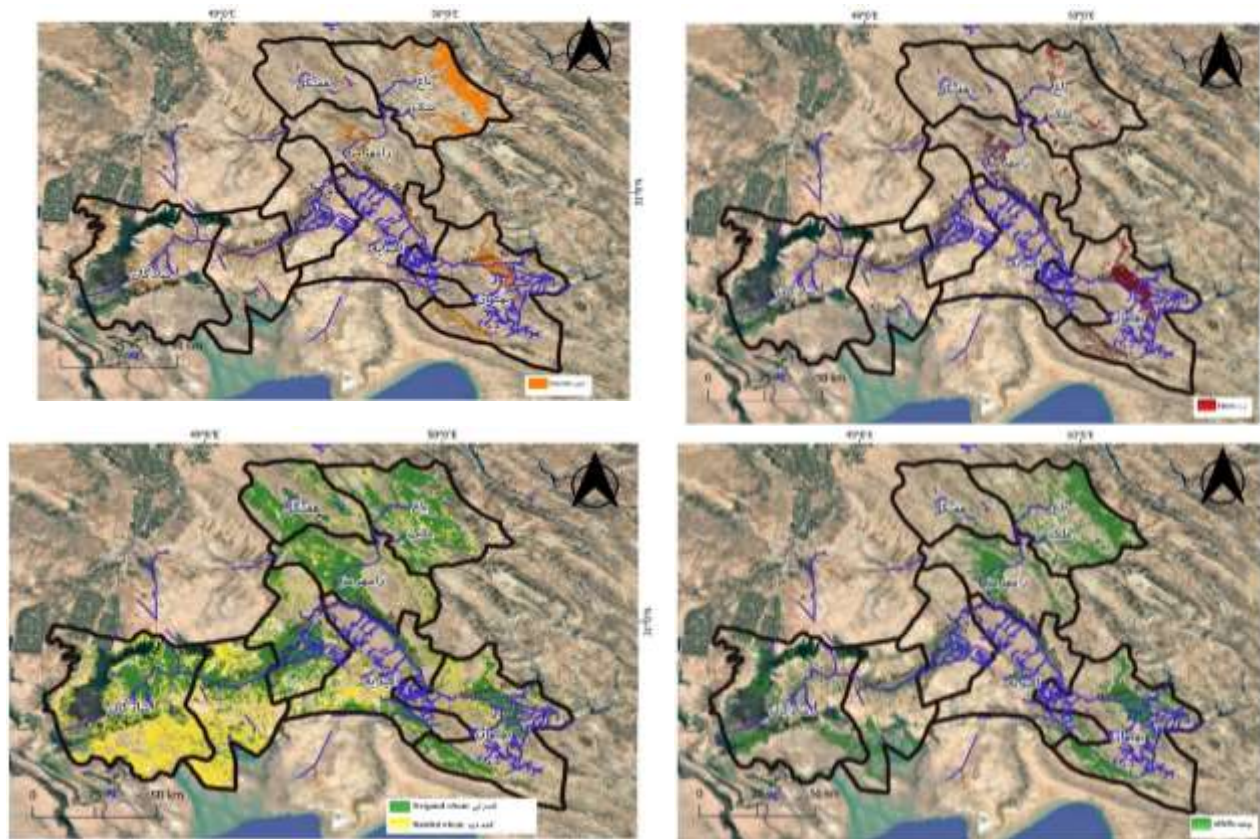
شکل ۴. ماتریس درهم ریختگی مدل SVM

در مجموع، تحلیل شاخص‌های طبقه‌ای نشان می‌دهد که عملکرد مدل در همه کلاس‌ها نسبتاً پایدار بوده اما از دقت بالا و تفکیک کامل برخوردار نیست. این موضوع به ویژه در کلاس‌هایی که از نظر ویژگی‌های طیفی هم‌پوشانی دارند (مانند کنگد، ذرت و یونجه) مشهود است. همچنین، تغییرات زمانی تصاویر سنتینل در بازه‌های رشد کنگد و ذرت احتمالاً موجب شده بخشی از تفاوت‌های فنولوژیکی در داده‌ها منعکس نشود، و همین امر باعث افت دقت در کلاس‌های دارای هم‌پوشانی طیفی شده است. این یافته‌ها با مطالعات اخیر نیز هم‌خوانی دارد. به عنوان مثال، سلطانی و بهمن‌آبادی (۱۴۰۴)، گزارش کردند که SVM در داده‌های چندطیفی با ابعاد بالا، در مقایسه با مدل‌های یادگیری مجموعه‌ای مانند RF و XGBoost، دقت پایین‌تری دارد، به ویژه زمانی که کلاس‌ها دارای شباهت طیفی بالایی باشند (Soltani).

همچنین (Kok et al., 2021) نشان دادند که اگرچه SVM به‌طور سنتی برای داده‌های با ابعاد بالا مناسب است، اما در کاربردهای سنجش از دور کشاورزی، به دلیل وجود شباهت طیفی میان محصولات مختلف، صحت آن کاهش می‌یابد. در شکل ۵، نقشه طبقه‌بندی محصولات کشاورزی با استفاده از روش SVM و جانمایی مزارع هر محصول به تفکیک ارائه شده است.

جدول ۵. نتایج شاخص‌های ارزیابی دقت به تفکیک کلاس‌های مورد بررسی

کلاس	Precision	Recall	F1-Score
ذرت	۰/۵۹	۰/۵۵	۰/۵۷
کنجد	۰/۴۵	۰/۵۴	۰/۴۹
یونجه	۰/۴۹	۰/۵۴	۰/۵۲
گندم	۰/۶۵	۰/۵۴	۰/۵۹
پهنه‌های آبی	۰/۷۱	۰/۵۴	۰/۶۱
خاک	۰/۵۱	۰/۵۴	۰/۵۳
چلیک	۰/۴۹	۰/۵۴	۰/۵۲



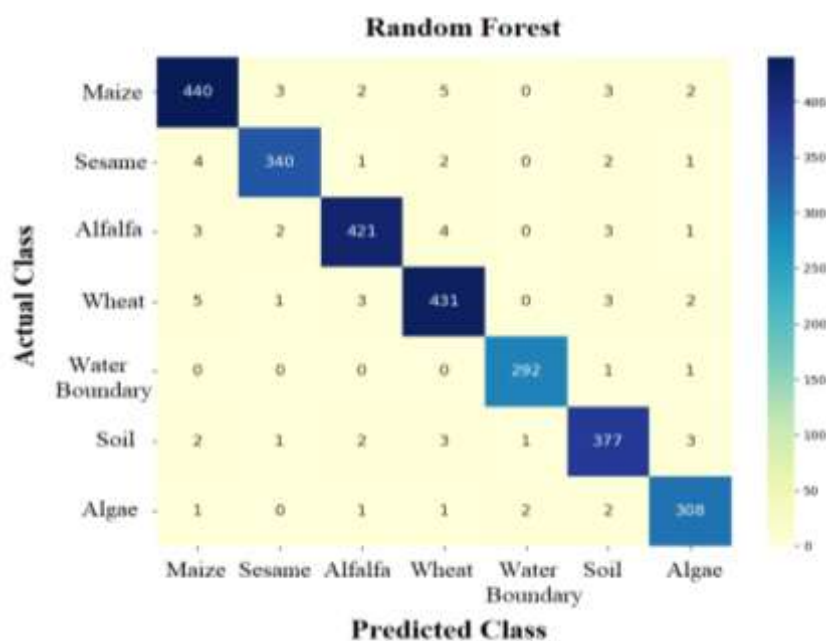
شکل ۵. نقشه طبقه‌بندی محصولات کشاورزی با استفاده از روش SVM

بررسی عملکرد جنگل تصادفی

ماتریس ابهام ارائه در شکل ۶ شده نشان می‌دهد که تعداد نمونه‌های درست طبقه‌بندی شده، برای تمامی کلاس‌ها نسبتاً بالا بوده و خطاهای طبقه‌بندی بین کلاس‌ها محدود است. این امر خود نشان‌دهنده توانایی مدل در شناسایی ویژگی‌های متمایز هر کلاس است. برای کلاس ذرت، ۴۴۰ نمونه به درستی شناسایی شده‌اند و تعداد کمی نمونه (جمعاً ۱۵ نمونه) به اشتباه به سایر کلاس‌ها تخصیص یافته است. براساس نتایج جدول ۶ این موضوع با شاخص‌های دقت، فراخوانی و F1-Score برابر با ۰/۹۶ مطابقت دارد و نشان‌دهنده توانایی بالای مدل در شناسایی ذرت است. خطاهای اندک عمدتاً به اشتباهات جزئی در تفکیک با کلاس‌های مشابه گیاهی (مانند گندم و یونجه) بازمی‌گردد، که ممکن است ناشی از شباهت‌های طیفی در برخی دوره‌های رشد باشد. کلاس کنجد نیز عملکرد بسیار بالایی دارد و از ۳۵۰ نمونه آموزشی، ۳۴۰ نمونه به درستی شناسایی شده‌اند Precision و Recall به ترتیب ۰/۹۸ و ۰/۹۷ هستند که منجر به F1-Score ۰/۹۷ شده است. خطاهای مشاهده شده شامل تخصیص محدود به کلاس‌های ذرت و یونجه است که می‌تواند ناشی از



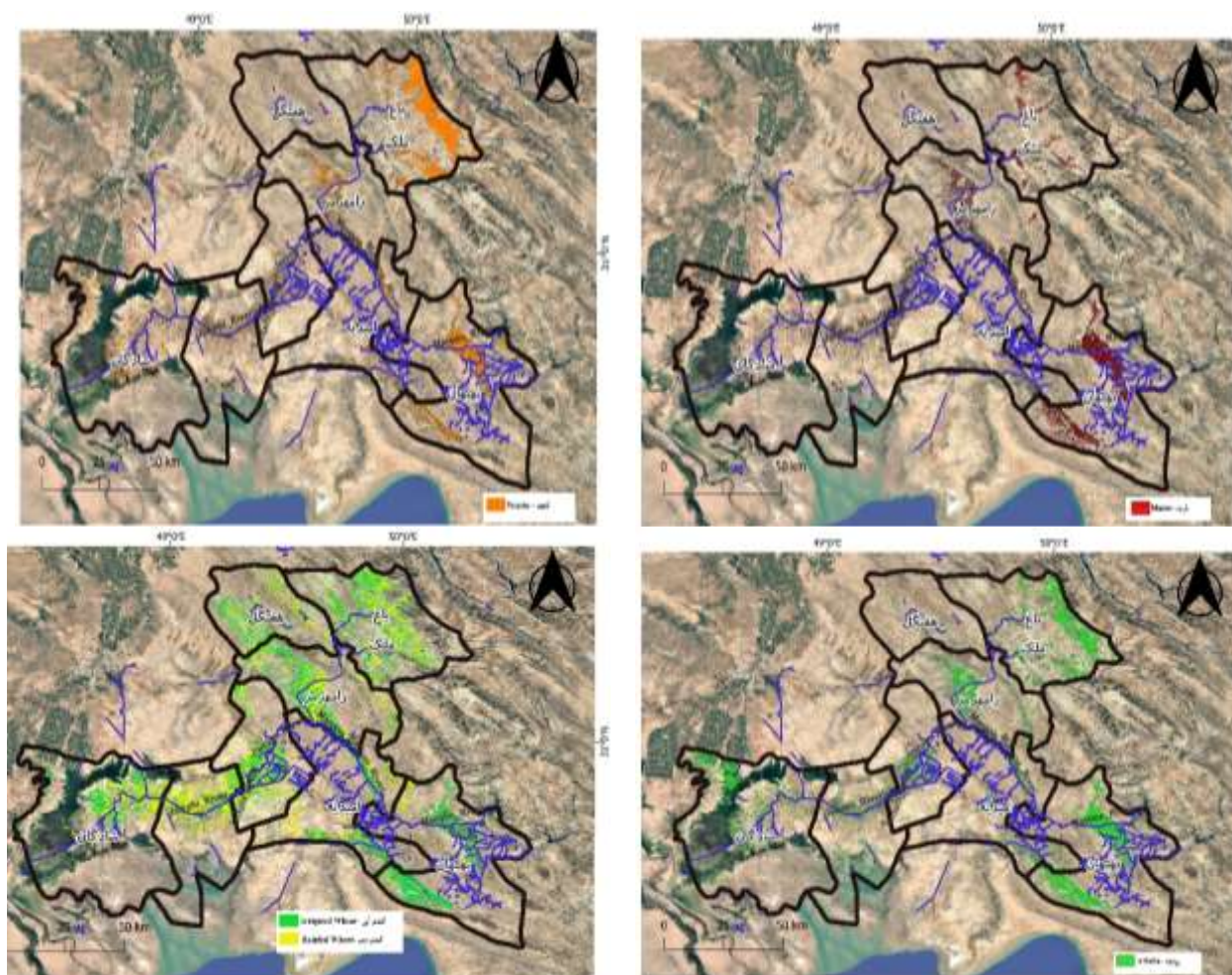
همپوشانی‌های جزئی در شاخص‌های گیاهی باشد. برای یونجه، ۴۲۱ نمونه از مجموع نمونه‌ها به درستی طبقه‌بندی شده‌اند و Precision، Recall و F1-Score در محدوده ۰/۹۷ است. این مقادیر نشان‌دهنده تفکیک مناسب یونجه از سایر گیاهان و خاک است، اگرچه خطاهای اندکی به کلاس‌های ذرت و گندم مشاهده می‌شود که قابل انتظار است، زیرا الگوهای رشد و ترکیب طیفی این گیاهان در برخی بازه‌های زمانی مشابه است. کلاس گندم با ۴۳۱ نمونه درست طبقه‌بندی شده، عملکرد بسیار خوب و قابل قبولی نشان می‌دهد Precision و Recall برابر ۰/۹۶ بوده و F1-Score، ۰/۹۷ است. برخی خطاهای جزئی به اشتباه در کلاس‌های ذرت و یونجه مشاهده می‌شود که به دلیل شباهت ویژگی‌های طیفی در مرحله‌ای از رشد ممکن است رخ داده باشد. کلاس خاک نیز با Precision و Recall و F1-Score، ۰/۹۶ عملکرد بالایی داشته است. خطاهای مشاهده شده عمدتاً به اشتباهات کوچک با کلاس‌های ذرت و گندم بازمی‌گردد که ناشی از ترکیب طیفی مشابه خاک باز در حاشیه مزارع است. در نهایت، کلاس جلبک با دقت بالای ۰/۹۶ در هر سه شاخص مورد بررسی، عملکرد بسیار مناسب و قابل توجهی دارد. خطاهای جزئی بین جلبک و مرز آب یا خاک مشاهده شده است، که می‌تواند به دلیل تغییرات ریز در سطح آب و تراکم جلبک‌ها باشد. به طور کلی، ماتریس ابهام و شاخص‌های دقت نشان می‌دهد که مدل جنگل تصادفی توانسته است تمامی کلاس‌ها را با دقت بالا و قابل قبول تفکیک کند. بالاتر بودن Precision و Recall برای کلاس‌های آب و جلبک بیانگر تمایز واضح ویژگی‌های فیزیکی و طیفی این کلاس‌ها از سایر پوشش‌هاست، در حالی که کمی کاهش شاخص‌ها برای کلاس‌های گیاهی می‌تواند ناشی از همپوشانی‌های فصلی و ویژگی‌های مشابه برخی گیاهان باشد. علاوه‌براین، نتایج بیانگر آن است که ترکیب داده‌های اپتیک و راداری، به‌ویژه ویژگی‌های بافتی سنتینل-۱ و شاخص‌های گیاهی سنتینل-۲، موجب افزایش چشمگیر قدرت تفکیک مدل جنگل تصادفی شده است. در شکل ۷، نقشه طبقه‌بندی محصولات کشاورزی با استفاده از روش Random Forest و جانمایی مزارع هر محصول به تفکیک ارائه شده است.



شکل ۶. ماتریس درهم ریختگی مدل جنگل تصادفی

جدول ۶. نتایج شاخص‌های ارزیابی دقت به تفکیک کلاس‌های مورد بررسی

F1-Score	Recall	Precision	کلاس
۰/۹۶	۰/۹۶	۰/۹۶	ذرت
۰/۹۷	۰/۹۷	۰/۹۸	کنجد
۰/۹۷	۰/۹۷	۰/۹۷	یونجه
۰/۹۶	۰/۹۶	۰/۹۶	گندم
۰/۹۹	۰/۹۹	۰/۹۹	پهنه‌های آبی
۰/۹۶	۰/۹۶	۰/۹۶	خاک
۰/۹۷	۰/۹۷	۰/۹۶	جلبک

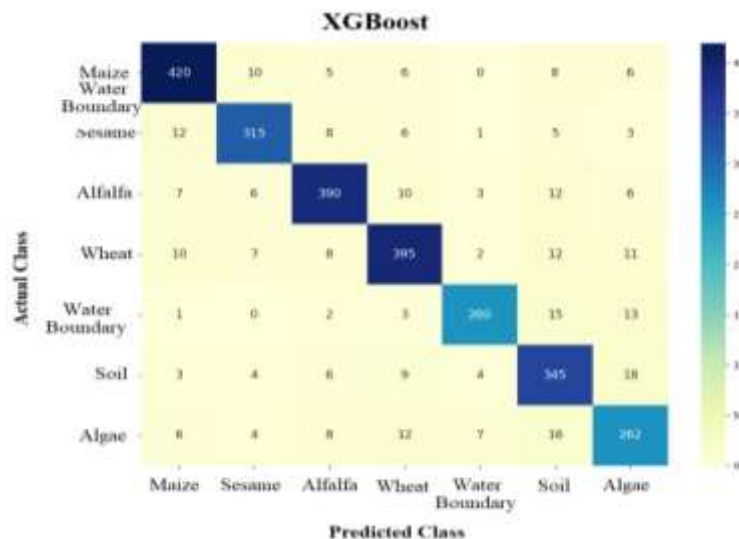


شکل ۷. نقشه طبقه‌بندی محصولات کشاورزی با استفاده از روش جنگل تصادفی

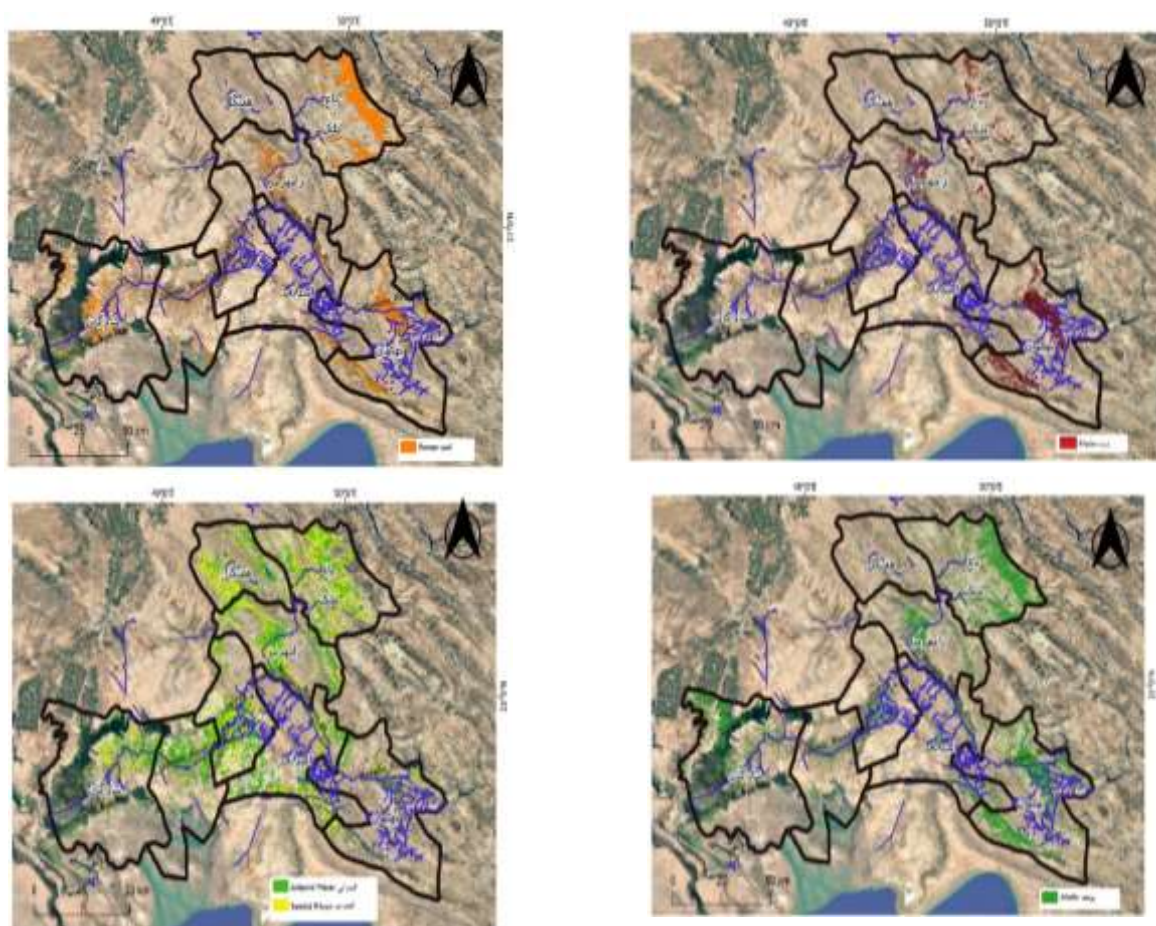
تحلیل عملکرد XGBoost

بر اساس ماتریس درهم‌ریختگی (شکل ۸)، مدل توانسته است کلاس‌های اصلی زراعی شامل ذرت، کنگد، یونجه و گندم را با دقت قابل قبولی تفکیک کند. براساس نتایج جدول ۷، کلاس ذرت با مقدار Precision و Recall بالای ۰/۹۰ نشان‌دهنده‌ی شناسایی صحیح اغلب پیکسل‌های مربوط به ذرت است؛ با این حال، تعداد اندکی از نمونه‌ها با کلاس‌های مجاور از جمله گندم و یونجه اشتباه گرفته شده‌اند. این همپوشانی می‌تواند ناشی از شباهت طیفی مراحل رشد این محصولات باشد که منجر به نزدیکی شاخص‌های طیفی و بافتی در تصاویر سنجنش‌ازدور می‌شود. کلاس کنگد نیز با F1-Score، ۰/۹ عملکرد مناسبی داشته است، اما در مقایسه با ذرت، اندکی افت در Recall مشاهده می‌شود. این موضوع احتمالاً به دلیل پوشش گیاهی تنک‌تر و بازتاب بالاتر خاک در مزارع کنگد است که موجب افزایش احتمال اشتباه با کلاس‌های خاک و گندم شده است. در مورد یونجه و گندم، مقادیر F1-Score به ترتیب ۰/۹۰ و ۰/۸۹ بوده است. هر دو کلاس عملکرد نسبتاً پایداری دارند، اما خطاهای متقابل میان آن‌ها نشان می‌دهد که تفکیک دقیق بین این دو محصول در برخی دوره‌های رشد دشوار است. در زمان‌های خاصی از فصل، هر دو محصول ممکن است شاخص‌های سبزیگی مشابهی مانند NDVI یا SAVI داشته باشند که موجب هم‌پوشانی طیفی در داده‌های سنتینل ۲ می‌شود. کلاس پهنه‌های آبی با بالاترین مقدار Precision و F1-Score در میان کلاس‌ها نشان‌دهنده‌ی قابلیت بالای مدل در تشخیص پهنه‌های آبی است. با این حال، کاهش اندک در مقدار Recall حاکی از آن است که بخش‌هایی از پهنه‌های آبی کم‌عمق یا دارای پوشش گیاهی آبی، به اشتباه در کلاس‌های دیگر قرار گرفته‌اند. در مقابل، کلاس‌های خاک و جلبک نسبت به سایر کلاس‌ها دقت پایین‌تری نشان می‌دهند. برای خاک، Precision، ۰/۸۳ و F1-Score، ۰/۸۶ بیانگر آن است که بخش‌هایی از خاک‌های مرطوب یا دارای پوشش گیاهی تنک به اشتباه در کلاس‌های زراعی ثبت شده‌اند. این امر به‌ویژه در مناطقی که بقایای گیاهی پس از برداشت باقی مانده یا رطوبت سطحی بالاست، بیشتر رخ می‌دهد. در مجموع، XGBoost در این مطالعه توانسته است با و توزیع متعادل خطاها، عملکردی قابل قبول در تفکیک کلاس‌های مختلف ارائه دهد. بیشترین دقت مربوط به کلاس‌های با

ویژگی‌های طیفی متمایز (مانند آب و ذرت) و کمترین دقت مربوط به کلاس‌های دارای همپوشانی طیفی (مانند جلبک و خاک) است. بنابراین، این مدل توانسته است تعادل مناسبی میان دقت کلی و پایداری در طبقه‌بندی کلاس‌های زراعی و غیرزراعی برقرار کرده و با انتخاب مناسب ویژگی‌ها و بهبود داده‌های ورودی، پایش دقیق و گسترده‌ی کاربری اراضی در مناطق کشاورزی را فراهم آورد. در عین حال، حساسیت بالاتر XGBoost به داده‌های نویزی یا نامتوازن ممکن است بخشی از افت دقت در کلاس‌های خاک و جلبک را توجیه کند. در شکل ۹، نقشه طبقه‌بندی محصولات کشاورزی با استفاده از روش Random Forest و جانمایی مزارع هر محصول به تفکیک ارائه شده است.



شکل ۸. ماتریس درهم ریختگی مدل XGBoost



شکل ۹. نقشه طبقه‌بندی محصولات کشاورزی با استفاده از روش XGBoost

جدول ۷. نتایج شاخص‌های ارزیابی دقت به تفکیک کلاس‌های مورد بررسی

کلاس	Precision	Recall	F1-Score
ذرت	۰/۹۱۵	۰/۹۲۳	۰/۹۱۹
کنجد	۰/۹۱۰	۰/۹	۰/۹۰۵
یونجه	۰/۹۱۳	۰/۸۹۹	۰/۹۰۶
گندم	۰/۸۹۶	۰/۸۸۸	۰/۸۹۲
پهنه‌های آبی	۰/۹۳	۰/۸۸۴	۰/۹۱۱
خاک	۰/۸۳۵	۰/۸۸۷	۰/۸۶۰
چلیک	۰/۸۲۱	۰/۸۳۲	۰/۸۲۶

مقایسه عملکرد سه مدل یادگیری ماشین در طبقه‌بندی تصاویر

براساس جدول ۸، نتایج این پژوهش نشان داد که سه مدل SVM، جنگل تصادفی و XGBoost در طبقه‌بندی محصولات زراعی و پوشش‌های زمینی بر اساس داده‌های ترکیبی سنتینل-۱، سنتینل-۲ و شاخص‌های گیاهی عملکرد متفاوتی از خود نشان دادند. الگوریتم SVM با F1-Score برابر با ۰/۵۵ و دقت کلی تنها ۰/۵۴ نشان می‌دهد که تقریباً نصف نمونه‌ها به درستی تفکیک نشده‌اند. این وضعیت در مقایسه با مطالعات پیشین که برای SVM در کارهای طبقه‌بندی پوشش/کاربری اراضی مقادیر دقت بسیار بالاتر گزارش کرده‌اند، قابل ملاحظه است. برای مثال، در مطالعه‌ای، دقت کلی برای SVM و RF به ترتیب حدود ۰/۸۶ و ۰/۸۳ گزارش شده است (Aduagna et al., 2022). این عملکرد پایین می‌تواند به دلیل همبستگی زیاد میان ویژگی‌ها باشد، چراکه SVM بر پایه فاصله‌ها در فضای ویژگی‌ها کار می‌کند و این عامل می‌تواند باعث شود که مرز تصمیم یا «هسته» آن نتواند به درستی تفکیک کلاس‌ها را انجام دهد (Dabija et al., 2021). از سوی دیگر، عملکرد بسیار بالای مدل جنگل تصادفی نشانه‌ای قوی از طراحی مناسب پژوهش است؛ یعنی به نظر می‌رسد داده‌های آموزشی از تنوع زمانی، مکانی و پوشش محصولی کافی برخوردار بوده‌اند، و فرآیند انتخاب ویژگی‌ها و مدل‌سازی به گونه‌ای بوده که توانسته است روابط پیچیده میان ویژگی‌ها و کلاس‌ها را به خوبی مدل نماید. مطالعات مروری نیز اشاره دارند که جنگل تصادفی به دلیل عملکرد خوب در برابر همپوشانی کلاس‌ها، توانایی مدیریت ویژگی‌های مختلف، و حساسیت کمتر نسبت به مقیاس‌دهی، یکی از گزینه‌های بسیار مناسب برای طبقه‌بندی پوشش و کاربری اراضی است.

در رابطه با مدل XGBoost، عملکرد آن با F1-Score و صحت کلی بالای ۸۵٪ نشانگر آن است که این مدل نیز گزینه‌ای بسیار مناسب بوده، اما با جنگل تصادفی فاصله داشته است. در مجموع، در شرایطی که داده‌ها چندمنبعی و دارای تنوع زمانی بالا هستند، جنگل تصادفی پایداری و دقت بالاتری دارد، در حالی که XGBoost در تنظیمات بهینه قادر است نتایج مشابه یا حتی برتر ارائه دهد.

مطالعات پیشین نشان داده‌اند که مدل RF در طبقه‌بندی محصولات زراعی با تصاویر سنتینل-۲ صحت بیشتری نسبت به SVM دارد؛ برای مثال، (Saini & Ghosh, 2018) صحت‌های کلی RF را حدود ۸۴/۲۲ درصد و ۸۱/۸۵ درصد برای SVM گزارش کرده‌اند. همچنین، (Tugac et al., 2024) نیز در مطالعه‌ای روی طبقه‌بندی محصولات زراعی با سنتینل-۱ و ۲ و لندست ۸ نشان دادند که RF عملکرد بهتری نسبت به SVM دارد. از سوی دیگر، استفاده از مدل XGBoost در ترکیب با داده‌های چندزمانه سنتینل-۱ و سنتینل-۲، منجر به بهبود قابل توجه در صحت کلی و ضریب کاپا شده است؛ به طوری که در مطالعه‌ای در جنوب شرقی فرانسه، صحت کلی به ۹۱/۰۹ درصد و کاپای برابر با ۰/۸۸ افزایش یافت؛ این بهبود نسبت به RF در همان تنظیمات قابل توجه بود.

در مجموع، مقایسه سه مدل نشان می‌دهد که RF پایدارترین و دقیق‌ترین روش برای طبقه‌بندی محصولات زراعی در این مطالعه بوده است، در حالی که XGBoost در پوشش‌های آبی برتری نسبی داشته و SVM ضعیف‌ترین عملکرد را از خود نشان داده است. بنابراین، برای کاربردهای کشاورزی و پایش کاربری اراضی، استفاده از RF یا ترکیب آن با مدل‌های پیشرفته‌تر مانند XGBoost می‌تواند بهترین نتیجه را فراهم آورد.

همچنین، در این پژوهش، مشاهده شد که هنگام کار در حوضه گسترده‌ای مانند مارون-جراحی، دو مدل SVM و XGBoost با خطای توقف زمانی^۱ مواجه شدند، در حالی که RF در شرایط مشابه با سرعت بالاتری اجرا شد. این تفاوت عملکرد را می‌توان با ساختار



محاسباتی و پیچیدگی الگوریتمی آنها توضیح داد. مدل SVM به‌ویژه در فرم غیرخطی نیازمند محاسبه و ذخیره ماتریس هسته هستند که به‌صورت مربعی در ابعاد $n \times n$ تعریف می‌شود. با افزایش n ، حافظه و زمان محاسباتی به‌سرعت افزایش می‌یابد. همچنین بسیاری از پیاده‌سازی‌های استاندارد SVM از روش‌هایی مانند SMO استفاده می‌کنند که پیچیدگی زمانی آن می‌تواند تا $O(n^3)$ باشد. افزون بر این، نیاز به حل مسائل بهینه‌سازی مربعی و اسکن‌های چندباره از داده‌ها، قابلیت کاربرد مدل را برای داده‌های بزرگ به‌شدت کاهش می‌دهد (Platt, 1998; Yu et al., 2003).

در سوی دیگر، XGBoost به‌عنوان یک مدل تقویتی پیشرفته برای داده‌های حجیم طراحی شده و دارای قابلیت‌هایی مانند موازی‌سازی، اجرای توزیع‌شده و محاسبات خارج از حافظه است. با این حال، در صورت افزایش عمق درخت‌ها، تعداد زیاد درختان یا تنظیمات سنگین، زمان اجرای آن به‌شدت افزایش می‌یابد. دلیل این امر، محاسبات مکرر گرادیان‌ها، منظم‌سازی و انتخاب شاخص‌های تقسیم^۴ در هر مرحله است که هزینه محاسباتی را بالا می‌برد. بنابراین حتی اگر XGBoost در حالت بهینه‌سازی شده کارآمد باشد، در داده‌های بسیار بزرگ و تنظیمات پیچیده می‌تواند با خطای توقف زمانی مواجه شود (Ramdani & Furqon, 2022). در مقابل، RF با ساخت مجموعه‌ای از درخت‌های تصمیم مستقل و اجرای موازی آنها، مدلی سبک‌تر و مقیاس‌پذیرتر است. هر درخت به‌طور مستقل و با نمونه‌برداری از داده آموزش می‌بیند و پیچیدگی زمانی آن به‌طور میانگین در حدود $O(n \log n \cdot d \cdot k)$ باقی می‌ماند، که در آن k تعداد درخت‌ها است. همین ساختار ساده‌تر و امکان اجرای موازی موجب می‌شود RF روی داده‌های بزرگ در زمان مناسب و بدون خطای توقف زمان، به نتیجه برسد (Ramdani & Furqon, 2022).

جدول ۸. مقایسه دقت و صحت مدل‌های یادگیری ماشین در طبقه‌بندی تصاویر

ML-Model	Overall Accuracy	Kappa Coefficient	Precision	Recall	F1-Score
SVM	۵۴/۹۲٪	۰/۴۷	۰/۵۶۱	۰/۵۴۹	۰/۵۵۲
Random Forest	۹۷/۲۸٪	۰/۹۷	۰/۹۷۴	۰/۹۷۴	۰/۹۷۴
XGBoost	۸۹٪	۰/۸۷	۰/۸۹	۰/۸۸۷	۰/۸۸۸

بررسی صحت مدل‌های یادگیری ماشین در برآورد سطح زیرکشت محصولات

جدول ۹، مقایسه‌ای بین سطح واقعی و سطح برآورد شده محصولات کشاورزی مختلف با استفاده از مدل‌های SVM، جنگل تصادفی و XGBoost را ارائه می‌دهد. نتایج نشان می‌دهد که هر یک از مدل‌ها رفتار متفاوتی در تخمین سطح زیرکشت داشته‌اند.

برای محصول گندم، مدل جنگل تصادفی، نزدیک‌ترین برآورد به مقدار واقعی سطح زیرکشت گندم را داشته است که اختلافی حدود ۵ درصد را نشان می‌دهد. این نتیجه بیانگر توانایی بالای مدل جنگل تصادفی در شناسایی گستره واقعی کشت گندم است و نشان می‌دهد که این مدل به خوبی قادر است الگوهای طیفی و زمانی مربوط به گندم را در مجموعه داده‌ها تفکیک کند. در مقابل، برآوردهای مدل SVM به وضوح کمتر از مقدار واقعی است و اختلاف بالایی دارد، و مدل XGB نیز برآورد میانگینی بین دو مدل دیگر ارائه کرده است. این اختلافات نشان می‌دهد که جنگل تصادفی در شناسایی گستره واقعی گندم نسبت به دو مدل دیگر عملکرد برتری دارد و SVM به دلیل حساسیت به همپوشانی کلاس‌ها و محدودیت در تفکیک الگوهای طیفی، سطح زیرکشت واقعی را کمتر از حد انتظار تخمین زده است. برای کنگد، مدل جنگل تصادفی، مقدار نزدیک‌تری (اختلاف ۸- درصد) نسبت به SVM و XGB (اختلاف ۳۱+ درصد) ارائه کرده است. این موضوع نشان می‌دهد که جنگل تصادفی قابلیت بالاتری در تفکیک الگوی طیفی و زمانی کنگد داشته است. این نتیجه نشان‌دهنده توانایی بالاتر جنگل تصادفی در تفکیک الگوی طیفی و زمانی کنگد است و تأکید می‌کند که مدل‌های مبتنی بر درخت تصمیم، به ویژه جنگل تصادفی، در تشخیص محصولات با ویژگی‌های طیفی مشابه و پراکندگی زمانی بهتر عمل می‌کنند.

در مورد یونجه، هر سه مدل سطح زیرکشت بیشتری نسبت به مقدار واقعی برآورد کرده‌اند، اما بیشترین خطا مربوط به SVM بوده است. این امر می‌تواند ناشی از همپوشانی طیفی یونجه با سایر محصولات یا حساسیت مدل SVM به پارامترها و ویژگی‌ها باشد. در این زمینه، جنگل تصادفی و XGB عملکرد نسبتاً بهتری ارائه داده‌اند، که بیانگر توانایی این مدل‌ها در مدیریت همپوشانی داده‌ها و مدل‌سازی

1 kernel matrix

2 Sequential Minimal Optimization

3 regularization

4 split

روابط غیرخطی میان ویژگی‌ها است.

در خصوص ذرت، هر سه مدل مقادیر نسبتاً مشابهی را تخمین زده‌اند و عملکرد مدل‌ها کمتر تفاوت قابل توجهی داشته است. در این بین، SVM کمترین اختلاف را با سطح واقعی داشته و پس از آن XGB و جنگل تصادفی قرار گرفته‌اند. این موضوع نشان می‌دهد که در شرایطی که ویژگی‌های طیفی محصول به‌وضوح مشخص هستند یا نمونه‌های آموزشی نمایندگی کافی از داده‌ها دارند، SVM می‌تواند عملکرد قابل قبولی ارائه دهد.

به طور کلی، تحلیل جدول ۹، نشان می‌دهد که مدل جنگل تصادفی بهترین عملکرد را در تخمین سطح زیرکشت گندم، یونجه و کنجد داشته است و توانایی آن در مدیریت همپوشانی کلاس‌ها و تفکیک الگوهای پیچیده، عملکرد بهتری نسبت به SVM و XGB ارائه داده است. در مقابل، SVM برای محصول ذرت عملکرد دقیق‌تری داشته است که نشان‌دهنده حساسیت این مدل به ویژگی‌های مشخص و نمونه‌های آموزشی مناسب است. این نتایج به وضوح نشان می‌دهند که انتخاب مدل باید بر اساس نوع محصول، ویژگی‌های طیفی و پراکندگی زمانی و مکانی داده‌ها انجام شود و نشانگر اهمیت ارزیابی مقایسه‌ای الگوریتم‌ها در مطالعات برآورد سطح زیرکشت است.

جدول ۹. مقایسه سطح زیر کشت برآورده شده براساس مدل‌های یادگیری ماشین و مقدار واقعی

Crop	Bias	سطح ریز کشت	Bias	سطح ریز کشت	Bias	سطح ریز کشت	Bias
Area Estimation (ha)		سطح ریز کشت گندم (ha)		سطح ریز کشت یونجه (ha)		سطح ریز کشت کنجد (ha)	
Actual Area (ha)		۱۳۵۶۵۸		۶۶۸۷		۷۰۱۲	
SVM	۰/۰۳	۳۴۹۲/۴۵	-۰/۱۳	۶۰۳۵/۵۵	۳/۵	۳۰۲۷۱/۳۳	-۰/۴۷
RF	-۰/۱۳	۲۸۸۲/۴۹	-۰/۰۸	۶۲۳۴/۶۰	۰/۰۷	۷۱۹۶/۲۴	۰/۰۵
XGB	۰/۱۹	۳۹۵۵/۶۳	۰/۳۱	۹۱۲۳/۵۴	۰/۱۹	۸۵۱۷/۰۱	-۰/۲۱

نتیجه‌گیری و پیشنهادها

این پژوهش با هدف مقایسه سه الگوریتم یادگیری ماشین شامل جنگل تصادفی، ماشین بردار پشتیبان و تقویت گرادیان شدید در تلفیق داده‌های ماهواره‌ای سنتینل-۱ و سنتینل-۲ به منظور طبقه‌بندی محصولات زراعی و برآورد سطح زیرکشت در حوضه مارون-جراحی انجام شد. یافته‌ها نشان داد که مدل جنگل تصادفی با صحت کلی ۹۷٪، ضریب کاپای ۰/۹۶ و F1-Score برابر ۰/۹۷، بالاترین دقت را در تفکیک محصولات کشاورزی (به‌ویژه گندم و کنجد) ارائه داد و توانست بیشترین انطباق را با داده‌های آماری واقعی ارائه دهد. مدل XGBoost نیز با صحت کلی ۸۹٪ عملکرد رضایت‌بخشی داشت و در تفکیک پوشش‌های آبی و مرز بین محصولات نتایج مناسبی نشان داد، اما تنظیم دقیق ابرپارامترها (مانند learning_rate و max_depth) می‌تواند دقت XGBoost را به جنگل تصادفی نزدیک کند.

در مقابل مدل SVM با صحت کلی ۵۴٪ به دلیل شباهت طیفی بین برخی محصولات و توزیع نامتوازن داده‌ها دقت پایین‌تری داشت. به‌طور کلی، یافته‌ها نشان داد که استفاده از شاخص‌های تکمیلی مانند F1-Score، Precision و Recall در کنار صحت کلی و ضریب کاپا، در ارزیابی واقع‌بینانه عملکرد مدل‌ها ضروری است. مدل‌های درختی مانند RF به دلیل پایداری بالا، مقاومت در برابر نویز و توانایی کار با داده‌های نامتوازن، گزینه‌ای قابل اعتماد برای پایش کشاورزی در مقیاس‌های ناحیه‌ای محسوب می‌شوند. همچنین، ترکیب خروجی مدل‌های RF و XGBoost در قالب رویکردهای Ensemble یا Voting می‌تواند در آینده موجب کاهش خطای برآورد سطح زیرکشت شود.

نتایج این پژوهش می‌تواند به‌عنوان مبنایی علمی برای سیاست‌گذاری در مدیریت بهینه منابع آب و زمین‌های کشاورزی در سطح حوضه‌های آبخیز مورد استفاده قرار گیرد. اطلاعات دقیق از سطح زیرکشت و الگوی زراعی به‌دست‌آمده از این مدل‌ها، امکان پایش تغییرات کاربری اراضی، تخصیص عادلانه منابع آبی، برنامه‌ریزی کشت منطقه‌ای و پیش‌بینی نیاز آبی محصولات را برای نهادهای تصمیم‌گیر (مانند وزارت جهاد کشاورزی و سازمان آب منطقه‌ای) فراهم می‌سازد. به این ترتیب، پژوهش حاضر علاوه بر جنبه‌های علمی، کاربرد مستقیم در مدیریت پایدار کشاورزی و سیاست‌گذاری امنیت غذایی در مناطق خشک و نیمه‌خشک کشور دارد.

به‌عنوان پیشنهاد، توسعه سامانه‌ای عملیاتی بر پایه داده‌های رایگان سنتینل و الگوریتم‌های یادگیری ماشین توصیه می‌شود تا پایش



سطح زیرکشت به صورت خودکار و تکرارپذیر انجام گیرد و ابزار تصمیم‌یار مؤثری برای مدیریت منابع و برنامه‌ریزی کشاورزی در مناطق خشک و نیمه‌خشک کشور فراهم گردد.

منابع

- باقری، نیکروز؛ سبزواری، علیرضا؛ رجبی‌پور، علی (۱۴۰۳). پایش سطح زیرکشت و طبقه‌بندی محصولات زراعی با استفاده از فناوری سنجش از دور ماهواره ای. *مهندسی زراعی*، ۴۷(۳)، ۳۳۷-۳۵۶. <https://doi.org/10.22055/agen.2024.46785.1723>
- خسروی، ایمان. (۲۰۲۴). تهیه نقشه نوع محصول کشاورزی از سری زمانی تصاویر لندست-۸ با استفاده از روش‌های یادگیری ماشین (مطالعه موردی: مرودشت استان فارس). *جغرافیا و برنامه ریزی محیطی*، ۳۵(۲)، ۴۵-۶۶. <https://doi.org/10.22108/gep.2024.138615.1601>
- رضائی اعتدالی، هادی؛ احمدی، مژگان. (۱۴۰۳). بررسی ارتباط بین شاخص‌های خشکسالی با عملکرد ذرت با استفاده از روش جنگل تصادفی (مطالعه موردی: شبکه آبیاری دشت قزوین). *نیوار*، ۴۸(۱۲۶-۱۲۷)، ۱۲۷-۱۳۷. <https://doi.org/10.30467/nivar.2024.467444.1299>
- سلطانی، مسعود؛ بهمن آبادی، بهاره. (۱۴۰۴). تحلیل عملکرد الگوریتم‌های یادگیری ماشین در طبقه‌بندی تصاویر چند زمانه در مدیریت کشاورزی (مطالعه موردی: دشت قزوین). *مدل سازی و مدیریت آب و خاک*، ۵(۳)، ۵۴-۷۳. <https://doi.org/10.22098/mmws.2025.16425.1532>
- شریفی پیچون، محمد؛ امیدوار، کامران؛ متذکر، کوثر. (۱۳۹۸). استفاده از روش تحلیل خوشه ای و رگرسیون چند متغیره در ارزیابی پتانسیل سنجی سیلاب با تاکید بر پارامترهای هیدروژئومورفولوژیکی (مورد مطالعه: حوضه آبخیز رودخانه مارون). *مخاطرات محیط طبیعی*، ۸(۲۱)، ۷۵-۹۲. <https://doi.org/10.22111/jneh.2018.22519.1336>
- طاهری دهکردی، علیرضا؛ ولدان زوج، محمدجواد. (۱۴۰۰). طبقه‌بندی کاربری اراضی کشاورزی با استفاده از تصاویر ماهواره‌ای سنتینل و شبکه عصبی پیچشی سه‌بعدی عمیق (مطالعه موردی: شهرکرد). *تحقیقات آب و خاک ایران*، ۵۲(۷)، ۱۹۴۱-۱۹۵۳. <https://doi.org/10.22059/ijswr.2021.320850.668956>
- قدسی، زینب؛ خیرخواه زرکش، میرمسعود؛ قرمزچشمه، باقر. (۱۳۹۹). مقایسه دقت روش‌های ماشین بردار پشتیبان و جنگل تصادفی در تهیه نقشه کاربری اراضی و محصولات زراعی، با استفاده از تصاویر چندزمانه سنتینل-۲. *نشریه سنجش از دور و GIS ایران*، ۱۲(۴)، ۷۳-۹۲. <https://doi.org/10.52547/gisj.12.4.73>
- نویدی، میرناصر؛ چترنور، منصور؛ جمشیدی، محمد؛ اخیانی، احمد. (۱۴۰۰). برآورد سطح کشت چند محصول انتخابی با استفاده از تصاویر چند زمانه سنتینل-۲ در دشت بسطام. *پژوهش‌های خاک*، ۳۵(۱)، ۴۱-۵۸. <https://doi.org/10.22092/ijsr.2021.353903.592>

REFERENCES

- Adugna, T., Xu, W., & Fan, J. (2022). Comparison of Random Forest and Support Vector Machine Classifiers for Regional Land Cover Mapping Using Coarse Resolution FY-3C Images. *Remote Sensing*, 14(3), 574. <https://doi.org/10.3390/rs14030574>
- Alejo, R., Garcia, V., Sotoca, J. M., Mollineda, R. A., & Sánchez, J. S. (2006). Improving the classification accuracy of RBF and MLP neural networks trained with imbalanced samples. *International Conference on Intelligent Data Engineering and Automated Learning*.
- Ali, A., Imran, M., & Khan, M. A. (2022). Evaluating Sentinel-2 red edge through hyperspectral profiles for monitoring LAI & chlorophyll content of Kinnow Mandarin orchards. <https://doi.org/10.1016/j.rsase.2022.100719>.
- Ashourloo, D., Nematollahi, H., Huete, A., Aghighi, H., Azadbakht, M., Shahrabi, H. S., & Goodarzashti, S. (2022). A new phenology-based method for mapping wheat and barley using time-series of Sentinel-2 images. *Remote Sensing of Environment*, 280, 113206. <https://doi.org/10.1016/j.rse.2022.113206>
- Aziz, G., Minallah, N., Saeed, A., Frnda, J., & Khan, W. (2024). Remote sensing based forest cover classification using machine learning. *Scientific Reports*, 14(1), 69. <https://doi.org/10.1038/s41598-023-50863-1>
- Bagheri, N., Sabzevari, A., & Rajabipour, A. (2024). Monitoring and classification of crops with satellite remote sensing technology. *Agricultural Engineering*, 47(3), 337-356. <https://doi.org/10.22055/agen.2024.46785.1723>. [In Persian]
- Baugh, W., & Groeneveld, D. (2006). Broadband vegetation index performance evaluated for a low-cover environment. *International Journal of Remote Sensing*, 27(21), 4715-4730.

- Breiman, L. (2001). Random Forests. *Machine Learning*, 45(1), 5-32. <https://doi.org/10.1023/A:1010933404324>
- Chen, T., & Guestrin, C. (2016). *XGBoost: A Scalable Tree Boosting System* Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, San Francisco, California, USA. <https://doi.org/10.1145/2939672.2939785>
- Clevers, J. G. P. W. (1989). Application of a weighted infrared-red vegetation index for estimating leaf Area Index by Correcting for Soil Moisture. *Remote sensing of environment*, 29(1), 25-37. [https://doi.org/https://doi.org/10.1016/0034-4257\(89\)90076-X](https://doi.org/https://doi.org/10.1016/0034-4257(89)90076-X)
- Dabija, A., Kluczek, M., Zagajewski, B., Raczko, E., Kycko, M., Al-Sulttani, A. H., Corbera, J. (2021). Comparison of Support Vector Machines and Random Forests for Corine Land Cover Mapping. *Remote Sensing*, 13(4), 777. <https://www.mdpi.com/2072-4292/13/4/777>
- Dong, J., Fu, Y., Wang, J., Tian, H., Fu, S., Niu, Z., . . . Yuan, W. (2020). Early season mapping of winter wheat in China based on Landsat and Sentinel images. *Earth System Science Data Discussions*, 2020, 1-26 .
- Feizizadeh, B., Darabi, S., Blaschke, T., & Lakes, T. (2022a). QADI as a New Method and Alternative to Kappa for Accuracy Assessment of Remote Sensing-Based Image Classification. *Sensors*, 22, (12)
- Feizizadeh, B., Darabi, S., Blaschke, T., & Lakes, T. (2022b). QADI as a New Method and Alternative to Kappa for Accuracy Assessment of Remote Sensing-Based Image Classification. *Sensors*, 22(12), 4506. <https://www.mdpi.com/1424-8220/22/12/4506>
- Foley, J. A., Ramankutty, N., Brauman, K. A., Cassidy, E. S., Gerber, J. S., Johnston, M., Zaks, D. P. M. (2011). Solutions for a cultivated planet. *Nature*, 478(7369), 337-342. <https://doi.org/10.1038/nature10452>
- Foody, G .M. (2023). Challenges in the real world use of classification accuracy metrics: From recall and precision to the Matthews correlation coefficient. *Plos one*, 18(10), e0291908 . <https://doi.org/10.1371/journal.pone.0291908>
- Gamal, E., Khder, G., Morsy, A., El-Sayed, M., Hashim, A., & Saleh, H. (2020). (Hyperspectral indices for discriminating plant diversity in Wadi AL-Afreet, Egypt. *Plant Archives*, 20(Supplement 2), 3361-3371 .
- George, R., Padalia, H., & Kushwaha, S. P. S. (2014). Forest tree species discrimination in western Himalaya using EO-1 Hyperion. *International Journal of Applied Earth Observation and Geoinformation*, 28, 140-149. <https://doi.org/https://doi.org/10.1016/j.jag.2013.11.011>
- Gessner, U., Hirner, A., Asam, S., Wenzl, M., & Kuenzer, C. (2025). Combining Machine Learning and Spatio-Temporal Filtering to Map Crop Types of Germany for 7 Years. *IEEE Geoscience and Remote Sensing Letters* .
- Ghanian, M. (2024). Exploring the food, energy, and water governance in South-West Iran. *Regional Science Policy & Practice*, 16(2), 12578. <https://doi.org/https://doi.org/10.1111/rsp3.12578>
- Ghods, Z., Kheirkhah Zarkesh, M. M., & Ghermezcheshmeh, B. (2021a). Comparison of Accuracy Between Support Vector Machine and Random Forest Classifiers for Land Use and Crop Mapping Using Multi-Temporal Sentinel-2 Images. *Iranian Journal of Remote Sensing & GIS*, 12(4), 73-92. <https://doi.org/10.52547/gisj.12.4.73> [In Persian]
- Goldblatt, R., You, W., Hanson, G., & Khandelwal, A. K. (2016). Detecting the Boundaries of Urban Areas in India: A Dataset for Pixel-Based Image Classification in Google Earth Engine. *Remote Sensing*, 8(8), 634. <https://www.mdpi.com/2072-4292/8/8/634>
- Haddadi, S. R., Hashemi, M., Peralta, R. C., & Soltani, M. (2025). Visible Image-Based Machine Learning for Identifying Abiotic Stress in Sugar Beet Crops. *Algorithms*, 18(11), 680. <https://doi.org/10.3390/a18110680>
- Hajirad, I., Mohammadi, S., & Dehghanisani, H. (2023). Determining the critical points of a basin from the point of view of water productivity and water consumption using the wapor database. *Environmental Sciences Proceedings*, 25(1), 86. <https://doi.org/10.3390/ECWS-7-14322>
- Halder, K., Srivastava, A. K., Zheng, W., Alsafadi, K., Zhao, G., Maerker, M., . . . Ewert, F. (2025). A robust and scalable crop mapping framework using advanced machine learning and optical and SAR imageries. *Smart Agricultural Technology*, 101354. <https://doi.org/https://doi.org/10.1016/j.atech.2025.101354>
- Hasan, S. F., Shareef, M. A., & Hassan, N. D. (2021). Speckle filtering impact on land use/land cover



- classification area using the combination of Sentinel-1A and Sentinel-2B (a case study of Kirkuk city, Iraq). *Arabian Journal of Geosciences*, 14(4), 276. <https://doi.org/10.1007/s12517-021-06494-9>
- Khosravi, I. (2024). Crop Mapping from Landsat-8 Images Time Series Using Machine-Learning Methods (Case Study: Marvdasht in Fars Province). *Geography and Environmental Planning*, 35(2), 45-66. <https://doi.org/10.22108/gep.2024.138615.1601> [In Persian]
- Khosravi, I., Razoumny, Y., Hatami Afkoeieh, J., & Alavipanah, S. K. (2021). Fully polarimetric synthetic aperture radar data classification using probabilistic and non-probabilistic kernel methods. *European Journal of Remote Sensing*, 54(1), 310-317. <https://doi.org/10.1080/22797254.2021.1924081>
- Kok, Z. H., Mohamed Shariff, A. R., Alfatni, M. S. M., & Khairunniza-Bejo, S. (2021). Support Vector Machine in Precision Agriculture: A review. *Computers and Electronics in Agriculture*, 191, 106546. <https://doi.org/10.1016/j.compag.2021.106546>
- Kumar, V., Lalotra, G. S., Sasikala, P., Rajput, D. S., Kaluri, R., Lakshmana, K., . . . Uddin, M. (2022). Addressing binary classification over class imbalanced clinical datasets using computationally intelligent techniques. *Healthcare*. <https://doi.org/10.3390/healthcare10071293>
- Lavreniuk, M., Kussul, N., Meretsky, M., Lukin, V., Abramov, S., & Rubel, O. (2017, 29 May-2 June 2017). Impact of SAR data filtering on crop classification accuracy. *IEEE First Ukraine Conference on Electrical and Computer Engineering (UKRCON)* ,
- Liang, S., Chen, Y., Liu, J., Ma, H., Li, W., Sucharitakul, P., Zhang, F. (2024). Mapping paddy rice cropping intensity and calendar in monsoon asia at 20 m resolution between 2018 and 2021 from multi-source satellite data using a sample-free algorithm. *Available at SSRN 4948283* .
- Liu, Y., Bachofen, C., Wittwer, R., Silva Duarte, G., Sun, Q., Klaus, V. H., & Buchmann, N. (2022). Using PhenoCams to track crop phenology and explain the effects of different cropping systems on yield. *Agricultural Systems*, 195, 103306. <https://doi.org/10.1016/j.agry.2021.103306>
- Marcinkowska, A., Zagajewski, B., Ochtyra, A., Jarocińska, A., Raczko, E., Kupková, L., Meuleman, K. (2014). Mapping vegetation communities of the Karkonosze National Park using APEX hyperspectral data and Support Vector Machines. *Miscellanea Geographica*, 18(2), 23-29. <https://doi.org/10.2478/mgrsd-2014-0007>
- Muhammad Zayyanul, A., & Projo, D. (2019). The effects of polynomial interpolation and resampling methods in geometric correction on the land-cover classification accuracy of Landsat-8 OLI imagery: a case study of Kulon Progo area, Yogyakarta. *Proc.SPIE* ,
- Navidi, M., Chatrenour, M., Jamshidi, M., & Akhyani, A. (2021). Estimating Cultivation Area of Some Selected Crops in Bastam Plain by Using Multi-Temporal Sentinel-2 Images. *Iranian Journal of Soil Research*, 35(1), 41-58. <https://doi.org/10.22092/ijsr.2021.353903.592> [In Persian]
- Nicolau, A. P., Dyson, K., Saah, D., & Clinton, N. (2024). Accuracy Assessment: Quantifying Classification Quality. In J. A. Cardille, M. A. Crowley, D. Saah, & N. E. Clinton (Eds.), *Cloud-Based Remote Sensing with Google Earth Engine: Fundamentals and Applications* (pp. 135-145). Springer International Publishing. https://doi.org/10.1007/978-3-031-26588-4_7
- Nițu, A., Florea, C., Ivanovici, M., & Racoviteanu, A. (2025). NDVI and Beyond: Vegetation Indices as Features for Crop Recognition and Segmentation in Hyperspectral Data. *Sensors*, 25(12), 3817. <https://www.mdpi.com/1424-8220/25/12/3817>
- Pal, M., & Mather, P. (2006). Some issues in the classification of DAIS hyperspectral data. *International Journal of Remote Sensing - INT J REMOTE SENS*, 27, 2895-2916. <https://doi.org/10.1080/01431160500185227>
- Phiri, D., Simwanda, M., Salekin, S., Nyirenda, V. R., Murayama, Y., & Ranagalage, M. (2020). Sentinel-2 Data for Land Cover/Use Mapping: A Review. *Remote Sensing*, 12(14), 2291. <https://www.mdpi.com/2072-4292/12/14/2291>
- Piaser, E., & Villa, P. (2023). Evaluating capabilities of machine learning algorithms for aquatic vegetation classification in temperate wetlands using multi-temporal Sentinel-2 data. *International Journal of Applied Earth Observation and Geoinformation*, 117, 103202. <https://doi.org/10.1016/j.jag.2023.103202>
- Platt, J. (1998). *Sequential Minimal Optimization: A Fast Algorithm for Training Support Vector Machines*. <https://www.microsoft.com/en-us/research/publication/sequential-minimal-optimization-a-fast-algorithm-for-training-support-vector-machines/>
- Raeisi, N., Moradi, S., & Scholz, M. (2022). Surface Water Resources Assessment and Planning with the QUAL2KW Model: A Case Study of the Maroon and Jarahi Basin (Iran). *Water*, 14(5), 705.

- <https://www.mdpi.com/2073-4441/14/5/705>
- Ramdani, F., & Furqon, M. (2022). The simplicity of XGBoost algorithm versus the complexity of Random Forest, Support Vector Machine, and Neural Networks algorithms in urban forest classification. *F1000Research*, 11, 1069. <https://doi.org/10.12688/f1000research.124604.1>
- Ramezani Etedali, H & ,Ahmadi, M. (2024). Investigating the Relationship Between Drought indices and Maize Yield Using Random Forest Method (Case Study: Qazvin Plain Irrigation Network). *Nivar*, 48(126-127), 127-137. <https://doi.org/10.30467/nivar.2024.467444.1299> [In Persian]
- Ray, S. (2019) Exploring machine learning classification algorithms for crop classification using Sentinel 2 data. *ISPRS - International Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences*, XLII-3/W6, 573-578. <https://doi.org/10.5194/isprs-archives-XLII-3-W6-573-2019>
- Reisi-Gahrouei, O., Homayouni, S., McNairn, H., Hosseini, M., & Safari, A. (2019). Crop biomass estimation using multi regression analysis and neural networks from multitemporal L-band polarimetric synthetic aperture radar data. *International Journal of Remote Sensing*, 40(17), 6822-6840. <https://doi.org/10.1080/01431161.2019.1594436>
- Reisi Gahrouei, O., Homayouni, S., & Safari, A. (2019). ESTIMATING CANOLA'S BIOPHYSICAL PARAMETERS FROM TEMPORAL, SPECTRAL, AND POLARIMETRIC IMAGERY USING MACHINE LEARNING APPROACHES. *The International Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences*, 42, 885-889 .
- Rokni, K., & Musa, T. A. (2019). Normalized difference vegetation change index: A technique for detecting vegetation changes using Landsat imagery. *Catena*, 178, 59-63 . <https://doi.org/10.1016/j.catena.2019.03.007>
- Rondeaux, G., Steven, M., & Baret, F. (1996). Optimization of soil-adjusted vegetation indices. *Remote sensing of environment*, 55(2), 95-107. [https://doi.org/10.1016/0034-4257\(95\)00186-7](https://doi.org/10.1016/0034-4257(95)00186-7)
- Rubel, O., Lukin, V., Rubel, A., & Egiazarian, K. (2021). Selection of Lee Filter Window Size Based on Despeckling Efficiency Prediction for Sentinel SAR Images. *Remote Sensing*, 13(10), 1887. <https://www.mdpi.com/2072-4292/13/10/1887>
- Sahin, E. K. (2020). Assessing the predictive capability of ensemble tree methods for landslide susceptibility mapping using XGBoost, gradient boosting machine, and random forest. *SN Applied Sciences*, 2(7), 1308. <https://doi.org/10.1007/s42452-020-3060-1>
- Saini, R., & Ghosh, S. K. (2018). CROP CLASSIFICATION ON SINGLE DATE SENTINEL-2 IMAGERY USING RANDOM FOREST AND SUPPOR VECTOR MACHINE. *Int. Arch. Photogramm. Remote Sens. Spatial Inf. Sci.*, XLII-5, 683-688. <https://doi.org/10.5194/isprs-archives-XLII-5-683-2018>
- Sarker, I. H. (2021). Machine learning: Algorithms, real-world applications and research directions. *SN computer science*, 2(3), 160 . <https://doi.org/10.1007/s42979-021-00592-x>
- Sharifi paichoon, M., Omidvar, K., & Motazaker, K. (2019). Assessment of flooding using cluster analysis and multivariable regression methods with emphasis on hydro geomorphological parameters (Case study: Maroon catchment). *Journal of Natural Environmental Hazards*, 8(21), 75-92. <https://doi.org/10.22111/jneh.2018.22519.1336> [In Persian]
- Şimşek, F. F. (2025). Comparison of Agricultural Crop Type Classifications with Different Machine Learning Algorithms (RF-SVM-ANN-XGBoost) by Generating Ground Truth Data from Farmer Declaration Parcels. *International Journal of Engineering and Geosciences*, 10(2), 207-220 . <https://doi.org/10.26833/ijeg.1552141>
- Soltani, M., & Bahmanabadi, B. (2025). Analysis of Machine Learning Algorithms' Performance in Multitemporal Image Classification for Agricultural Management (Case Study: Qazvin Plain). *Water and Soil Management and Modelling*, 5(3), 54-73. <https://doi.org/10.22098/mmws.2025.16425.1532> [In Persian]
- Sun, L., & Schulz, K. (2015). The Improvement of Land Cover Classification by Thermal Remote Sensing. *Remote Sensing*, 7(7), 8368-8390. <https://www.mdpi.com/2072-4292/7/7/8368>
- Taheri Dehkordi, A., & Valadan Zoej, M. J. (2021). Classification of Croplands Using Sentinel-2 Satellite Images and a Novel Deep 3D Convolutional Neural Network (Case Study: Shahrekord). *Iranian Journal of Soil and Water Research*, 52(7), 1941-1953. <https://doi.org/10.22059/ijswr.2021.320850.668956> [In Persian]
- Trinder, J., & Salah, M. (2012). Aerial images and LiDAR data fusion for disaster change detection. *ISPRS Annals of the Photogrammetry, Remote Sensing and Spatial Information Sciences*, 1, 227-232 .



- Tugac, M., Şimşek, F. h. F., & Torunlar, H. (2024). Classification of Agricultural Crops with Random Forest and Support Vector Machine Algorithms Using Sentinel-2 and Landsat-8 Images. *International Journal of Environment and Geoinformatics*, 11, 106-118. <https://doi.org/10.30897/ijgeo.1479116>
- Vapnik, V. (2013). *The nature of statistical learning theory*. Springer science & business media .
- Yadav, S., Sharma, S., & Chaudhary, P. (2024). Review on Advancement in Machine Learning and Deep Learning Techniques for Crop Classification. In (pp. 593-610). <https://doi.org/10.56155/978-81-955020-7-3-53>
- Yang, H., Ming, B., Nie, C., Xue, B., Xin, J., Lu, X., . . . Li, S. (2022). Maize Canopy and Leaf Chlorophyll Content Assessment from Leaf Spectral Reflectance: Estimation and Uncertainty Analysis across Growth Stages and Vertical Distribution. *Remote Sensing*, 14(9), 2115. <https://www.mdpi.com/2072-4292/14/9/2115>
- Yu, H., Yang, J., & Han, J. (2003). *Classifying Large Data Sets Using SVM with Hierarchical Clusters*. <https://doi.org/10.1145/956750.956786>
- Zaka, M. M., & Samat, A. (2024). Advances in Remote Sensing and Machine Learning Methods for Invasive Plants Study: A Comprehensive Review. *Remote Sensing*, 16(20), 3781. <https://www.mdpi.com/2072-4292/16/20/3781>
- Zalaki Badili, N., Sayyad, G., Hemadi, K., Akhavan, S., & Abdi, A. (2013). Simulation of Runoff on Murun Dam Watershed (Idanak) using by. *Agricultural Engineering*, 35(2), 25-36. https://agrieng.scu.ac.ir/article_10005_c2d28605122e959337e2fccd68c9b385.pdf
- Zheng, Y., Dong, W., ZhipingYang, Lu, Y., Zhang, X., Dong, Y., & Sun, F. (2024). A new attention-based deep metric model for crop type mapping in complex agricultural landscapes using multisource remote sensing data. *International Journal of Applied Earth Observation and Geoinformation*, 134, 104204. <https://doi.org/https://doi.org/10.1016/j.jag.2024.104204>
- zoratipour, a., shahpari far, g & yusefi, A. (2025). Modeling vegetation changes in Zagros forests based on machine learning algorithms and satellite images. *Iranian Journal of Soil and Water Research*, -. <https://doi.org/10.22059/ijswr.2025.394576.669934>