

Comparing Machine Learning Algorithms for Land Cover Classification in Urban Areas

Rasoul Lotfi¹ , Hamed Faroqi^{2✉}

1.Civil Engineering Dept., Faculty of Engineering, University of Kurdistan, Sanandaj, Iran. E-mail: rasoulotfi2023@gmail.com

2.Corresponding Author, Civil Engineering Dept., Faculty of Engineering, University of Kurdistan, Sanandaj, Iran. E-mail: h.faroqi@uok.ac.ir

Article Info

Article type:
Research Article

Article history:

Received

2025-12-10

Received in revised form

2026-02-24

Accepted

2026-03-09

Available online

2026-03-29

Keywords:

Land Use and Land Cover, Changes (LULC), Sentinel-2, Random Forest, Machine Learning, Satellite Image Classification)

ABSTRACT

Accurate knowledge of Land Use and Land Cover (LULC) changes play a vital role in sustainable natural resource management, urban planning, and environmental monitoring. This study aims to evaluate and compare the performance of machine learning algorithms and statistical methods for land use classification in the Naysar area, a suburb of Sanandaj city.

In this study, Sentinel-2 satellite data, along with a set of input features including spectral bands, spectral indices (NDVI, MNDWI, SAVI, and NDBI), and topographic information, were utilized. Three classification approaches—Random Forest (RF), Maximum Likelihood Classification (MLC), and unsupervised K-means clustering—were implemented. Training and validation samples were collected through visual interpretation of high-resolution imagery and field surveys, and were used for model development and accuracy assessment.

The results showed that the Random Forest algorithm achieved the best performance, with an overall accuracy of 98% and a Kappa coefficient of 0.95, while the Maximum Likelihood method and K-means algorithm achieved overall accuracies of 95% and 91%, respectively.

The findings indicate that the integration of spectral and topographic features with the Random Forest algorithm provides an efficient and reliable approach for accurate land use classification and environmental change monitoring in urban and peri-urban areas.

Cite this article: Lotfi, Rasoul., Faroqi, Hamed. (2026). Comparing Machine Learning Algorithms for Land Cover Classification in Urban Areas. *Advanced Modeling in Civil Engineering*, 3(1),36-54.

DOI:10.22126/amcen.2026.13598.1071



© The Author(s).

DOI:10.22126/amcen.2026.13598.1071

Publisher: Razi University

Introduction

Accurate knowledge of Land Use and Land Cover (LULC) changes is essential for urban planning, natural resource management, and environmental monitoring. Rapid population growth and urban expansion have significantly accelerated LULC changes, leading to degradation of natural ecosystems. Traditional field-based mapping methods, despite their accuracy, are limited in spatial and temporal coverage, while remote sensing data provide an efficient alternative for large-scale and continuous monitoring. Satellite imagery such as Sentinel-2 has been widely used due to its high spatial and spectral resolution.

Recent studies highlight the effectiveness of machine learning algorithms, particularly Random Forest (RF), in improving classification accuracy compared to traditional statistical methods such as Maximum Likelihood Classification (MLC). In contrast, unsupervised approaches like K-means are mainly used for exploratory analysis and generally provide lower accuracy. Despite these advancements, challenges such as spectral similarity between classes and limited integration of multi-source data remain.

To address these gaps, this study proposes a comprehensive comparative framework integrating spectral bands, spectral indices (NDVI, MNDWI, SAVI, NDBI), and topographic data. The performance of RF, MLC, and K-means is evaluated using multiple accuracy metrics, providing a robust and generalized approach for LULC classification in complex urban environments.

Method

The study area is Naysar, a peri-urban region located near Sanandaj in western Iran, characterized by rapid urban expansion and significant land use changes in recent years. Sentinel-2 satellite imagery from the Copernicus program was used as the primary data source, providing multispectral information with spatial resolutions of 10–60 m.

Preprocessing steps included cloud and shadow masking, temporal filtering, and band integration using the Google Earth Engine (GEE) platform. A set of spectral indices, including NDVI (Normalized Difference Vegetation Index), MNDWI (Modified Normalized Difference Water Index), SAVI (Soil Adjusted Vegetation Index), and NDBI (Normalized Difference Built-up Index), along with topographic slope, were extracted as input features.

Reference data were collected through visual interpretation of high-resolution imagery and field surveys. A total of 100 samples across four classes (built-up, vegetation, bare land, and water bodies) were selected and randomly divided into training (70%) and validation (30%) datasets.

Three classification approaches representing different paradigms were implemented: Random Forest (RF) as a machine learning method, Maximum Likelihood Classification (MLC) as a statistical baseline, and K-means as an unsupervised clustering algorithm. All models were trained and evaluated under identical conditions to ensure a fair comparison.

Accuracy assessment was conducted using a confusion matrix and multiple evaluation metrics, including Overall Accuracy (OA), Kappa coefficient, Producer's Accuracy (PA), and User's Accuracy (UA), providing a comprehensive evaluation framework for model performance.

Results

The performance of Maximum Likelihood (ML), Random Forest (RF), and K-means algorithms was evaluated using overall accuracy, Kappa coefficient, producer's accuracy, user's accuracy, and class area statistics. ML achieved an overall accuracy of 95% and a Kappa coefficient of 0.92, indicating acceptable performance; however, misclassification was observed mainly in water and vegetation classes due to spectral overlap. RF provided the highest accuracy with 98% overall accuracy and a Kappa of 0.95, demonstrating superior classification performance across all land cover classes. This superiority is attributed to its ability to model complex nonlinear relationships and integrate multiple spectral and auxiliary features. In contrast, K-means, as an unsupervised method, achieved lower performance (91% overall accuracy and 0.87 Kappa) and showed limitations in separating spectrally similar classes, particularly water and vegetation. Results confirm that supervised algorithms significantly outperform unsupervised approaches in heterogeneous urban environments. Among all methods, RF produced the most stable class distribution and the lowest classification error. Variations in estimated class areas were more evident in K-means, especially for vegetation and water classes. Overall, RF demonstrated the highest robustness and reliability for LULC classification in the study area.

Conclusions

In this study, the performance of three classification algorithms, namely Maximum Likelihood (ML), Random Forest (RF), and K-means, was evaluated for land use/land cover classification using Sentinel-2 imagery. The results showed that RF achieved the highest overall accuracy and Kappa coefficient, indicating superior classification performance compared to the other methods.

Overall, RF demonstrated the most reliable and stable performance across all land cover classes due to its ability to model complex nonlinear relationships and handle multi-dimensional feature spaces. ML provided acceptable results but showed lower accuracy, mainly due to limitations related to statistical assumptions and spectral overlap between classes. K-means, as an unsupervised method, achieved the lowest accuracy, although it was able to capture the general spatial pattern of land cover distribution.

The findings confirm that supervised machine learning methods outperform unsupervised approaches in complex urban environments. In addition, the results highlight the importance of using multi-source features, including spectral indices and auxiliary data, to improve classification performance.

Overall, this study demonstrates that RF is a robust and effective method for LULC mapping in urban and peri-urban areas using Sentinel-2 data.

Author Contributions

All authors participated in writing and revising the article.

Conflict of Interest

Authors declared no conflict of interest.



مقایسه الگوریتم های یادگیری ماشین در طبقه بندی پوشش مناطق شهری

رسول لطفی^۱، حامد فاروقی^۲

۱. گروه مهندسی عمران، دانشکده فنی مهندسی، دانشگاه کردستان، سنندج، ایران. رایانامه: rasoulotfi2023@gmail.com

۲. نویسنده مسئول، گروه مهندسی عمران، دانشکده فنی مهندسی، دانشگاه کردستان، سنندج، ایران. رایانامه: h.farوقي@uok.ac.ir

اطلاعات مقاله	چکیده
نوع مقاله:	آگاهی دقیق از تغییرات کاربری و پوشش زمین (LULC) نقش مهمی در مدیریت پایدار منابع طبیعی، برنامه ریزی شهری و پایش محیط زیست دارد. این پژوهش به ارزیابی و مقایسه عملکرد الگوریتم های یادگیری ماشین و روش های آماری در طبقه بندی کاربری زمین در منطقه نایسر (حومه شهر سنندج) می پردازد. در این مطالعه، داده های ماهواره ای Sentinel-2 و مجموعه ای از ویژگی ها شامل باندهای طیفی، شاخص های طیفی (NDVI, MNDWI, SAVI, NDBI) و اطلاعات توپوگرافی به عنوان ورودی استفاده شدند و سه رویکرد طبقه بندی شامل جنگل تصادفی (RF)، روش حداکثر احتمال (MLC) و خوشه بندی بدون ناظر K-means به کار گرفته شد. نمونه های آموزشی و ارزیابی از طریق تفسیر بصری تصاویر با وضوح بالا و بازدید میدانی جمع آوری و برای آموزش و اعتبارسنجی مدل ها استفاده شدند. نتایج نشان داد الگوریتم جنگل تصادفی با دقت کلی ۹۸٪ و ضریب کاپا ۰/۹۵ بهترین عملکرد را ارائه می دهد، در حالی که روش حداکثر احتمال دقت کلی ۹۵٪ و الگوریتم K-means دقت ۹۱٪ داشت. این یافته ها نشان می دهند که ترکیب ویژگی های طیفی و توپوگرافی با الگوریتم جنگل تصادفی رویکردی کارآمد و قابل اعتماد برای طبقه بندی دقیق کاربری زمین و پایش تغییرات محیطی در مناطق شهری و پیرامونی است.
مقاله پژوهشی	
تاریخ دریافت:	
۱۴۰۴/۰۹/۱۹	
تاریخ بازنگری:	
۱۴۰۴/۱۲/۰۵	
تاریخ پذیرش:	
۱۴۰۴/۱۲/۱۸	
تاریخ انتشار:	
۱۴۰۵/۰۱/۰۹	
کلیدواژه ها:	تغییرات کاربری و پوشش زمین، سنتینل-۲، جنگل تصادفی، یادگیری ماشین
تغییرات کاربری و	
پوشش زمین،	
سنتینل-۲،	
جنگل تصادفی،	
یادگیری ماشین	

استناد: لطفی، رسول، فاروقی حامد. (۱۴۰۵). مقایسه الگوریتم های یادگیری ماشین در طبقه بندی پوشش مناطق شهری. مجله مدل سازی پیشرفته در

DOI:10.22126/amcen.2026.13598.1071

مهندسی عمران، ۳(۱)، ۵۴-۳۶.



© نویسندگان.

ناشر: دانشگاه رازی.

۱. مقدمه

دانش و آگاهی درباره تغییرات کاربری زمین و پوشش زمین (LULC) نقش حیاتی در برنامه‌ریزی شهری، مدیریت منابع طبیعی و پایش تغییرات محیطی دارد [۱،۲]. افزایش جمعیت و گسترش شهرنشینی از مهم‌ترین عوامل مؤثر بر تغییرات LULC محسوب می‌شوند که منجر به تغییر کاربری اراضی، تخریب پوشش‌های طبیعی و کاهش منابع اکولوژیکی شده‌اند [۳].

در مطالعات اولیه، نقشه‌برداری کاربری زمین عمدتاً بر پایه برداشت‌های میدانی انجام می‌شد که علی‌رغم دقت بالا، به دلیل هزینه زیاد و محدودیت مکانی و زمانی، برای مناطق وسیع کارایی محدودی داشت [۴]. در همین راستا، داده‌های سنجنش از دور به‌عنوان جایگزینی کارآمد، امکان پایش گسترده و مداوم تغییرات سطح زمین را فراهم کرده‌اند [۵،۶].

تصاویر ماهواره‌ای نظیر SPOT، Landsat و Sentinel به‌طور گسترده در مطالعات LULC مورد استفاده قرار گرفته‌اند و نقش مهمی در استخراج اطلاعات مکانی-زمانی ایفا کرده‌اند [۷،۸]. علاوه بر این، توسعه پلتفرم‌های پردازش ابری مانند Google Earth Engine (GEE) امکان تحلیل سری‌های زمانی حجیم را تسهیل کرده است [۱۰،۱۱].

مطالعات اخیر نشان داده‌اند که داده‌های Sentinel-2 به دلیل وضوح مکانی و طیفی بالا، به‌طور گسترده در ترکیب با الگوریتم‌های یادگیری ماشین برای طبقه‌بندی کاربری و پوشش زمین مورد استفاده قرار گرفته‌اند. به‌عنوان مثال، الگوریتم جنگل تصادفی (Random Forest) در بسیاری از مطالعات به‌عنوان یکی از دقیق‌ترین روش‌های طبقه‌بندی بر روی داده‌های Sentinel-2 گزارش شده است و عملکرد آن نسبت به روش‌های کلاسیک بهبود قابل توجهی نشان داده است. همچنین ترکیب شاخص‌های طیفی با الگوریتم‌های یادگیری ماشین مانند SVM و XGBoost نیز در مطالعات اخیر برای افزایش دقت طبقه‌بندی مورد استفاده قرار گرفته است. علاوه بر این، استفاده از رویکردهای بدون نظارت مانند K-means نیز به‌عنوان ابزار اولیه تحلیل ساختار داده در برخی

پژوهش‌ها گزارش شده است که بیشتر در مراحل اکتشافی داده کاربرد دارد [۱،۴،۷،۱۵،۱۷].

در کنار پیشرفت داده‌های سنجنش از دور، توسعه الگوریتم‌های یادگیری ماشین موجب ارتقای قابل توجه دقت طبقه‌بندی شده است. الگوریتم‌های نظارت‌شده مانند جنگل تصادفی (RF)، ماشین بردار پشتیبان (SVM^۲) و شبکه‌های عصبی مصنوعی (ANN^۳) به‌طور گسترده در طبقه‌بندی LULC مورد استفاده قرار گرفته‌اند و عملکرد مطلوبی در داده‌های پیچیده داشته‌اند [۱۲،۱۳]. در مقابل، روش‌های بدون نظارت مانند الگوریتم میانگین‌گیری K (K-means) در شرایط نبود داده آموزشی کاربرد دارند، اما معمولاً دقت پایین‌تری نسبت به روش‌های نظارت‌شده ارائه می‌دهند [۱۴].

مطالعات اخیر نشان می‌دهند که ترکیب داده‌های چندمنبعی و استفاده از الگوریتم‌های پیشرفته مانند XGBoost^۴ و یادگیری عمیق می‌تواند دقت طبقه‌بندی را به‌طور قابل توجهی افزایش دهد [۱۵،۱۷]. همچنین مشخص شده است که عملکرد مدل‌ها به شدت تحت تأثیر کیفیت داده‌های ورودی، انتخاب ویژگی‌ها و تنظیم پارامترهای الگوریتم است [۱۶].

در میان الگوریتم‌های مختلف، جنگل تصادفی به دلیل توانایی در مدیریت داده‌های پیچیده و جلوگیری از بیش‌برازش، به‌عنوان یکی از دقیق‌ترین روش‌ها در مطالعات LULC شناخته شده است [۱۸-۲۱]. مطالعات مقایسه‌ای نشان داده‌اند که RF نسبت به روش‌های کلاسیک عملکرد بهتری دارد و دقت بالاتری در طبقه‌بندی ارائه می‌دهد [۲۲،۲۵].

با وجود پیشرفت‌های موجود، چالش‌هایی نظیر همپوشانی طیفی بین کلاس‌ها، انتخاب بهینه ویژگی‌ها و محدودیت در مقایسه همزمان الگوریتم‌های مختلف همچنان باقی است.

با وجود حجم قابل توجهی از پژوهش‌های انجام‌شده در حوزه طبقه‌بندی کاربری و پوشش زمین (Land Use/Land Cover - LULC)، مرور نظام‌مند ادبیات نشان می‌دهد که بخش عمده‌ای از مطالعات پیشین یا بر توسعه و به‌کارگیری یک الگوریتم منفرد متمرکز بوده‌اند، یا در بهترین حالت به مقایسه محدود میان چند

^۲ Support Vector Machine

^۳ Artificial Neural Network

^۴ Extreme Gradient Boosting

^۱ Google Earth Engine (GEE)

دور مورد استفاده قرار می‌گیرد. الگوریتم حداکثر احتمال (Maximum Likelihood) در این پژوهش به‌عنوان یک روش کلاسیک مبتنی بر مبانی آماری و توزیع احتمالاتی داده‌ها انتخاب شده است. هدف از به‌کارگیری این الگوریتم، ایجاد یک خط مبنا (Baseline) در چارچوب مقایسه‌ای مطالعه است تا امکان ارز همچنین الگوریتم K-means به‌عنوان یکی از پرکاربردترین روش‌های بدون نظارتیابی میزان بهبود حاصل از روش‌های نوین یادگیری ماشین نسبت به رویکردهای سنتی فراهم گردد.

این الگوریتم با فرض توزیع نرمال چندمتغیره کلاس‌ها، یکی از روش‌های استاندارد و پرکاربرد در طبقه‌بندی داده‌های سنجش از دور محسوب می‌شود و در بسیاری از مطالعات کلاسیک به‌عنوان روش مرجع مورد استفاده قرار گرفته است. بنابراین نقش آن در این پژوهش، نه به‌عنوان یک روش ضعیف، بلکه به‌عنوان یک مدل مرجع برای مقایسه منصفانه عملکرد الگوریتم‌های پیشرفته‌تر تعریف شده است. همچنین الگوریتم K-means به‌عنوان یک روش بدون نظارت مبتنی بر ساختار طبیعی داده‌ها انتخاب شده است تا توانایی روش‌های غیرنظارتی در تفکیک کلاس‌ها بدون استفاده از داده‌های آموزشی مورد ارزیابی قرار گیرد [۳۴،۳۵].

از منظر داده و ویژگی‌ها، این پژوهش با بهره‌گیری همزمان از چند منبع اطلاعاتی شامل باندهای طیفی Sentinel-2، مجموعه‌ای از شاخص‌های طیفی (NDVI، NDBI، SAVI، MNDWI) و یک پارامتر کلیدی توپوگرافی (شیب)، یک فضای ویژگی چندمنبعی و غنی ایجاد کرده است. این هم‌جوشی سطح ویژگی (Feature-level Fusion) موجب افزایش قدرت تفکیک بین کلاس‌ها، کاهش اثر همپوشانی طیفی میان کاربری‌های مشابه و بهبود پایداری نتایج در شرایط پیچیده شهری شده است [۳۶].

در بخش ارزیابی، عملکرد مدل‌ها تنها بر اساس دقت کلی مورد بررسی قرار نگرفته است، بلکه یک چارچوب ارزیابی چندمعیاره و مکمل شامل دقت کلی (Overall Accuracy)، ضریب کاپا (Kappa Coefficient)، دقت تولیدکننده (Producer's Accuracy) و دقت کاربر (User's Accuracy) به‌کار گرفته شده است. این رویکرد امکان تحلیل دقیق‌تر رفتار هر الگوریتم در سطح کلاس‌های مختلف و همچنین ارزیابی واقعی‌تر قابلیت تعمیم‌پذیری آن‌ها را فراهم می‌سازد.

روش مشابه اکتفا کرده‌اند. در بسیاری از این پژوهش‌ها، ساختار تحلیل به‌گونه‌ای بوده است که امکان ارزیابی جامع و چندبعدی عملکرد الگوریتم‌ها فراهم نشده است. افزون بر این، در اغلب مطالعات موجود، ارزیابی دقت مدل‌ها عمدتاً به یک معیار کلی مانند Overall Accuracy محدود شده و سایر شاخص‌های مهم نظیر ضریب کاپا یا دقت‌های کلاس‌محور (Producer's/User's Accuracy) کمتر مورد توجه قرار گرفته‌اند. همچنین، بسیاری از این پژوهش‌ها صرفاً بر داده‌های طیفی تکیه داشته‌اند و از ترکیب سایر منابع اطلاعاتی مانند شاخص‌های طیفی و داده‌های توپوگرافی استفاده نکرده‌اند؛ موضوعی که در محیط‌های شهری پیچیده می‌تواند منجر به کاهش قدرت تفکیک‌پذیری کلاس‌ها و افت قابلیت تعمیم نتایج شود [۹،۲۳].

در نتیجه، مهم‌ترین خلأ در ادبیات علمی این حوزه، نبود یک چارچوب یکپارچه، استاندارد و قابل تعمیم برای مقایسه همزمان الگوریتم‌های مختلف یادگیری ماشین در کنار بهره‌گیری از داده‌های چندمنبعی است. به بیان دقیق‌تر، مطالعات موجود کمتر به طراحی ساختاری پرداخته‌اند که بتواند به‌صورت همزمان: (۱) الگوریتم‌های نظارت‌شده و بدون نظارت را در یک چارچوب واحد مقایسه کند، (۲) از ترکیب ویژگی‌های چندمنبعی شامل باندهای طیفی، شاخص‌های طیفی و اطلاعات توپوگرافی بهره‌برد، و (۳) عملکرد مدل‌ها را با استفاده از مجموعه‌ای از معیارهای چندگانه و مکمل ارزیابی نماید. این کمبود ساختاری به‌ویژه در مطالعات مربوط به محیط‌های شهری پیچیده، منجر به محدودیت در تحلیل واقعی رفتار الگوریتم‌ها و کاهش قابلیت تعمیم نتایج شده است [۳۴،۳۶،۳۷].

در پاسخ به این شکاف پژوهشی، مطالعه حاضر یک چارچوب مقایسه‌ای جامع و چندلایه برای ارزیابی عملکرد الگوریتم‌های طبقه‌بندی LULC ارائه می‌دهد. در این چارچوب، سه الگوریتم با ماهیت‌های متفاوت انتخاب شده‌اند تا امکان تحلیل تطبیقی میان رویکردهای مختلف فراهم گردد. الگوریتم جنگل تصادفی (Random Forest) به‌عنوان یک روش یادگیری جمعی مبتنی بر درخت‌های تصمیم انتخاب شده است که به دلیل توانایی بالا در مدل‌سازی روابط غیرخطی، مقاومت در برابر بیش‌برازش و عملکرد مناسب در داده‌های پیچیده، به‌طور گسترده در مطالعات سنجش از

۲-۲. داده‌های ماهواره‌ای سنتینل-۲

در این تحقیق، از داده‌های ماهواره سنتینل-۲ متعلق به برنامه کوپرنیک اتحادیه اروپا استفاده شد. این ماهواره با ارائه تصاویری با قدرت تفکیک مکانی ۱۰، ۲۰ و ۶۰ متر و قدرت تفکیک رادیومتریکی ۱۲ بیتی، امکان تحلیل دقیق پوشش‌ها و کاربری‌های مختلف سطح زمین را فراهم می‌کند. سنتینل-۲ دارای ۱۳ باند طیفی در محدوده مرئی، فروسرخ نزدیک (NIR) و فروسرخ کوتاه‌موج (SWIR) است که این تنوع طیفی توانایی بالایی در تفکیک کلاس‌هایی مانند گیاهان، منابع آبی، خاک و مناطق شهری ایجاد می‌کند. استفاده از داده‌های سنتینل-۲ به دلایل مختلفی انتخاب شد که از مهم‌ترین آن‌ها می‌توان به دقت مکانی بالای ۱۰ متر، پوشش طیفی گسترده با ۱۳ باند و همچنین دسترسی آزاد و رایگان به داده‌ها اشاره کرد. این ویژگی‌ها موجب شده است که سنتینل-۲ به یکی از منابع اصلی و کارآمد برای پایش تغییرات کاربری اراضی و تحلیل‌های چندزمانه در مناطق شهری و طبیعی تبدیل شود.

۲-۲-۱. پیش پردازش‌های لازم

پیش‌پردازش داده‌های سنجنش از دور یکی از مراحل کلیدی در بهبود کیفیت داده‌ها و افزایش دقت نتایج طبقه‌بندی محسوب می‌شود. هدف از این مرحله، آماده‌سازی تصاویر ماهواره‌ای برای استخراج ویژگی‌ها و اجرای الگوریتم‌های طبقه‌بندی است.

در این پژوهش، داده‌های مورد استفاده از نوع بازتاب سطحی^۵ (Surface Reflectance) بوده و از طریق پلتفرم Google Earth Engine (GEE) در دسترس قرار گرفته‌اند. این داده‌ها از قبل تحت فرآیندهای اولیه تصحیح رادیومتریکی و اتمسفری قرار گرفته‌اند؛ بنابراین نیازی به انجام مجدد این اصلاحات در این مطالعه وجود نداشته است.

با این حال، برای افزایش دقت تحلیل‌ها، مراحل تکمیلی پیش‌پردازش انجام شده است که شامل موارد زیر می‌باشد:

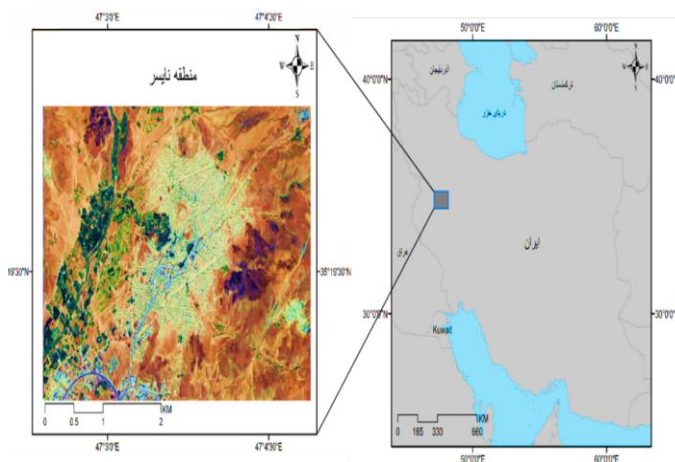
- انتخاب بازه زمانی مناسب برای کاهش اثر تغییرات فصلی
- ماسک کردن ابرها و سایه‌ها جهت حذف پیکسل‌های نامعتبر
- یکپارچه‌سازی باندهای مورد نیاز برای تحلیل

در مجموع، نوآوری این پژوهش نه تنها در مقایسه همزمان سه خانواده متفاوت از الگوریتم‌های طبقه‌بندی، بلکه در توسعه یک چارچوب یکپارچه مبتنی بر هم‌جوشی داده‌های چندمنبعی و ارزیابی چندمعیاره است. این چارچوب می‌تواند به عنوان یک الگوی قابل تعمیم برای مطالعات آینده در حوزه پایش و تحلیل تغییرات کاربری و پوشش زمین در محیط‌های پیچیده شهری مورد استفاده قرار گیرد.

۲. مواد و روش‌ها

۲-۱. منطقه مورد مطالعه

نایسر، یکی از نواحی حومه‌ای سنندج، مرکز استان کردستان در غرب ایران، به شمار می‌رود. این منطقه به خاطر افزایش سریع شهرنشینی و تغییرات بارز در پوشش و کاربری زمین در سال‌های اخیر، به‌ویژه تغییر اراضی کشاورزی و طبیعی به مناطق مسکونی و صنعتی، توجه ویژه‌ای را به خود جلب کرده است. توسعه سریع نایسر به دلیل نزدیکی آن به سنندج و قیمت پایین‌تر زمین و مسکن در مقایسه با مرکز شهر، این ناحیه را به نمونه‌ای برجسته از تحولات شهری در ایران تبدیل کرده است. نایسر در موقعیت جغرافیایی با طول‌های ۶۸۷۷۶۱.۰۸ متر شرقی و ۳۹۱۰۷۴۱.۸۱ متر شمالی واقع شده و بخشی از دامنه‌های زاگرس میانی محسوب می‌شود.



شکل ۱. نمای کلی از منطقه نایسر

⁵ Surface Reflectance

داده‌ها برای آموزش مدل‌ها (Training Data) و ۳۰ درصد برای ارزیابی صحت نتایج (Validation Data) مورد استفاده قرار گرفت [۳۰]. این تقسیم‌بندی استاندارد، امکان ارزیابی قابل اعتماد عملکرد الگوریتم‌ها را فراهم کرده و مبنای محاسبه شاخص‌های دقت شامل دقت کلی، دقت تولیدکننده، دقت کاربر و ضریب کاپا قرار گرفته است.

در مجموع، طراحی مناسب داده‌های مرجع در این پژوهش، نقش مهمی در افزایش قابلیت اعتماد نتایج طبقه‌بندی و مقایسه دقیق عملکرد الگوریتم‌های مختلف ایفا کرده است. تقسیم داده‌ها به مجموعه‌های آموزشی و آزمون به صورت مستقل انجام شده و مجموعه آزمون صرفاً برای ارزیابی نهایی مدل‌ها استفاده شده است.

۴-۲. روش‌شناسی (Methodology)

این بخش به تشریح مراحل اصلی روش تحقیق شامل استخراج ویژگی‌ها، اعمال الگوریتم‌های طبقه‌بندی و ارزیابی دقت نتایج اختصاص دارد. فرآیند کلی به گونه‌ای طراحی شده است که ابتدا ویژگی‌های مؤثر از داده‌های سنجش از دور استخراج شده، سپس طبقه‌بندی کاربری و پوشش زمین با استفاده از الگوریتم‌های مختلف انجام شده و در نهایت دقت نتایج مورد ارزیابی قرار گیرد. به منظور افزایش شفافیت و قابلیت تکرارپذیری نتایج، چارچوب روش‌شناسی این پژوهش به صورت ساختاریافته و در شرایط یکسان برای تمامی الگوریتم‌ها طراحی شده است. در این راستا، از داده‌های بازتاب سطحی Sentinel-2 در محیط Google Earth Engine استفاده شده و مراحل پیش‌پردازش شامل انتخاب بازه زمانی مناسب، ماسک‌کردن ابرها و سایه‌ها و یکپارچه‌سازی باندهای مورد نیاز به صورت استاندارد انجام شده است.

نمونه‌های مرجع شامل چهار کلاس اصلی کاربری زمین با بهره‌گیری از ترکیب تفسیر بصری تصاویر با قدرت تفکیک بالا و بازدیدهای میدانی تهیه شده‌اند. این داده‌ها به صورت تصادفی به دو بخش آموزش (۷۰ درصد) و اعتبارسنجی (۳۰ درصد) تقسیم شده‌اند تا از بروز سوگیری در فرآیند آموزش جلوگیری شده و امکان ارزیابی دقیق عملکرد مدل‌ها فراهم گردد.

- آماده‌سازی داده‌ها برای استخراج شاخص‌های طیفی و اجرای مدل‌های طبقه‌بندی
- اجرای این مراحل باعث افزایش کیفیت داده‌های ورودی و کاهش خطاهای احتمالی در فرآیند طبقه‌بندی نهایی شده و بستر مناسبی برای تحلیل دقیق‌تر کاربری اراضی فراهم کرده است.

۳-۲. داده‌های مرجع

داده‌های مرجع (Reference Data) یکی از مهم‌ترین اجزای مطالعات طبقه‌بندی در سنجش از دور محسوب می‌شوند و نقش اساسی در فرآیند آموزش، تنظیم و ارزیابی دقت مدل‌های طبقه‌بندی دارند [۳۰]. کیفیت و دقت این داده‌ها به طور مستقیم بر عملکرد الگوریتم‌های یادگیری ماشین تأثیرگذار بوده و به عنوان معیار اصلی برای اعتبارسنجی نتایج نهایی مورد استفاده قرار می‌گیرند. از این رو، انتخاب و تهیه نمونه‌های مرجع مناسب، یکی از مراحل کلیدی در تحلیل‌های مبتنی بر تصاویر ماهواره‌ای به شمار می‌رود.

در این پژوهش، داده‌های مرجع با بهره‌گیری از ترکیبی از تفسیر بصری تصاویر ماهواره‌ای با قدرت تفکیک مکانی مناسب و همچنین بازدیدهای میدانی تهیه شده‌اند [۳۰]. استفاده همزمان از این دو رویکرد موجب افزایش دقت در تعیین کلاس‌های کاربری زمین و کاهش خطاهای احتمالی ناشی از تفسیر صرف تصاویر شده است. فرآیند نمونه‌برداری به گونه‌ای طراحی شد که تمامی کلاس‌های اصلی موجود در منطقه مطالعه به طور مناسب پوشش داده شوند. در مجموع، تعداد ۱۰۰ نمونه مرجع برای چهار کلاس اصلی شامل مناطق شهری (Built-up)، پوشش گیاهی (Vegetation)، اراضی خاکی یا بایر (Bare land) و پهنه‌های آبی (Water bodies) تعریف گردید. انتخاب این کلاس‌ها بر اساس ویژگی‌های غالب کاربری زمین در منطقه مورد مطالعه و همچنین اهداف تحقیق صورت گرفته است. این نمونه‌ها به صورت تصادفی و با رعایت توزیع مناسب مکانی انتخاب شده‌اند تا از بروز سوگیری در فرآیند آموزش مدل جلوگیری شود.

به منظور ارزیابی دقیق عملکرد الگوریتم‌های طبقه‌بندی و کاهش احتمال بیش‌برازش (Overfitting)، داده‌های مرجع به دو بخش آموزشی و اعتبارسنجی تقسیم شدند. به این صورت که ۷۰ درصد از

شاخص نرمال‌شده تفاوت ساختارهای شهر (NDBI⁹) برای شناسایی و تفکیک مناطق ساخته‌شده از سایر کاربری‌ها [۳۴].

شیب (Slope) به‌عنوان یک پارامتر توپوگرافی مؤثر در تحلیل الگوهای مکانی [۳۵].

این شاخص‌ها با استفاده از باندهای طیفی تصاویر سنتینل-۲ و داده‌های کمکی استخراج شده‌اند و به‌عنوان مجموعه ویژگی‌های ورودی برای الگوریتم‌های طبقه‌بندی مورد استفاده قرار گرفته‌اند. ترکیب این ویژگی‌ها موجب افزایش قدرت تفکیک بین کلاس‌های مختلف و بهبود عملکرد مدل‌های یادگیری ماشین در محیط‌های پیچیده شهری و طبیعی شده است.

جدول ۱. روابط مربوط به شاخص‌های استفاده‌شده در روش پیشنهادی

شاخص	رابطه
Modified Normalized Difference Water Index (MNDWI)	$(\text{Green-SWIR}) / (\text{Green}+\text{SWIR})$
Normalized Difference Vegetation Index (NDVI)	$(\text{NIR-Red}) / (\text{NIR}+\text{Red})$
Normalized Difference Built-up Index (NDBI)	$\text{SWIR1- NIR} / \text{SWIR1}+\text{NIR}$
Soil Adjusted Vegetation Index (SAVI)	$((\text{NIR} - \text{Red}) / (\text{NIR} + \text{Red} + 0.5))$

همچنین، به‌منظور انجام یک مقایسه منصفانه، تمامی الگوریتم‌های طبقه‌بندی بر روی یک مجموعه داده مشترک و با شرایط یکسان اجرا شده‌اند. این رویکرد امکان تحلیل دقیق‌تر تفاوت عملکرد الگوریتم‌ها را فراهم کرده و از تأثیر عوامل جانبی بر نتایج جلوگیری می‌کند. لازم به ذکر است تمامی مراحل آموزش و ارزیابی مدل‌ها با رعایت اصل استقلال داده‌ها انجام شده است. به این صورت که داده‌های مرجع ابتدا به دو بخش آموزش (۷۰ درصد) و اعتبارسنجی (۳۰ درصد) تقسیم شدند و تمامی فرآیندهای آموزش مدل، تنظیم پارامترها و یادگیری الگوریتم‌ها صرفاً بر روی مجموعه داده‌های آموزشی انجام گرفته است. داده‌های آزمون به‌طور کامل از فرآیند آموزش و تنظیم مدل‌ها جدا نگه داشته شده و تنها برای ارزیابی نهایی عملکرد مدل‌ها مورد استفاده قرار گرفته‌اند. همچنین هیچ‌گونه بهینه‌سازی یا انتخاب ویژگی مبتنی بر داده‌های آزمون صورت نگرفته است تا از بروز هرگونه بیش‌برازش (Overfitting) جلوگیری شده و نتایج حاصل از طبقه‌بندی از نظر آماری معتبر و قابل تعمیم باشند.

۲-۴-۱. استخراج ویژگی

در تحلیل‌های مبتنی بر تصاویر سنجش از دور، استخراج ویژگی‌های مناسب نقش بسیار مهمی در بهبود دقت طبقه‌بندی و افزایش قابلیت تفکیک کلاس‌های مختلف کاربری زمین دارد. در این پژوهش، مجموعه‌ای از شاخص‌های طیفی و یک متغیر توپوگرافی به‌عنوان ورودی مدل‌های طبقه‌بندی مورد استفاده قرار گرفته‌اند.

این ویژگی‌ها شامل شاخص‌های کاربرد سنجش از دور هستند که هرکدام اطلاعات خاصی از ویژگی‌های سطح زمین ارائه می‌دهند: شاخص نرمال‌شده تفاوت پوشش گیاهی (NDVI⁶)، برای شناسایی و ارزیابی تراکم و سلامت پوشش گیاهی [۳۱]. شاخص نرمال‌شده تفاوت آب اصلاح‌شده (MNDWI⁷) برای تفکیک پهنه‌های آبی از سایر سطوح [۳۲]. شاخص پوشش گیاهی تعدیل‌شده خاک (SAVI⁸) برای کاهش اثر بازتاب خاک در مناطق کم‌پوشش گیاهی [۳۳].

⁶ Normalized Difference Vegetation Index

⁷ Modified Normalized Difference Water Index

⁸ Soil Adjusted Vegetation Index

⁹ Normalized Difference Built-up Index

متغیر ورودی بدون نیاز به فرض‌های آماری سخت‌گیرانه، آن را به گزینه‌ای مناسب برای تحلیل داده‌های حاصل از هم‌جوشی چندمنبعی تبدیل کرده است [۳۴، ۳۵].

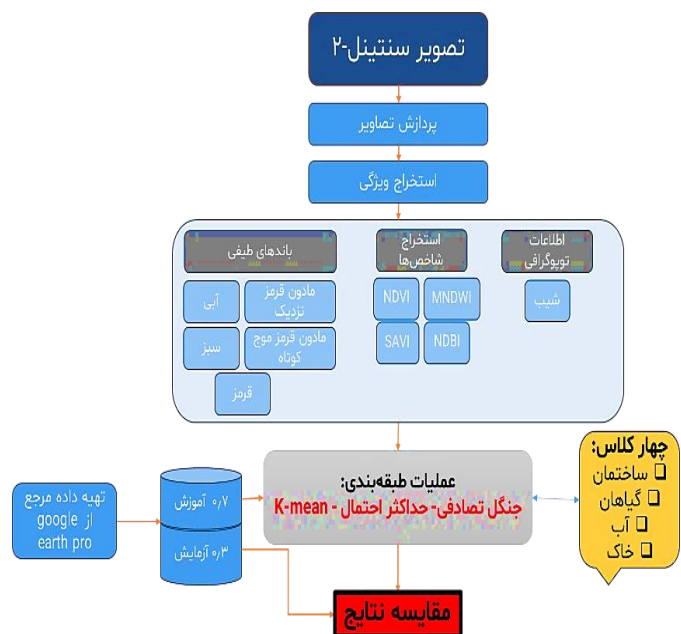
در مقابل، الگوریتم حداکثر احتمال (Maximum Likelihood) به‌عنوان یک روش کلاسیک آماری مبتنی بر فرض توزیع نرمال چندمتغیره انتخاب شده است. استفاده از این الگوریتم به‌عنوان یک خط مبنا (Baseline) امکان مقایسه مستقیم بین رویکردهای سنتی و روش‌های نوین یادگیری ماشین را فراهم می‌سازد و کمک می‌کند میزان بهبود حاصل از به‌کارگیری مدل‌های پیشرفته‌تر به‌صورت کمی ارزیابی گردد. علاوه بر این، به‌کارگیری یک روش آماری کلاسیک، درک بهتری از محدودیت‌ها و مزایای مدل‌های مبتنی بر فرضیات توزیعی در برابر روش‌های داده‌محور فراهم می‌کند.

همچنین، الگوریتم K-means به‌عنوان یکی از پرکاربردترین روش‌های بدون نظارت مبتنی بر خوشه‌بندی داده‌ها در این مطالعه مورد استفاده قرار گرفته است. این الگوریتم بدون نیاز به داده‌های آموزشی، ساختار طبیعی داده‌ها را بر اساس شباهت‌های درونی آن‌ها استخراج می‌کند. گنجانیدن این روش در چارچوب مقایسه‌ای تحقیق، امکان بررسی تفاوت عملکرد میان رویکردهای وابسته به داده‌های آموزشی و روش‌های مستقل از آن‌ها را فراهم ساخته و به تحلیل جامع‌تری از قابلیت‌های طبقه‌بندی در شرایط محدودیت داده‌های مرجع کمک می‌کند.

در مجموع، انتخاب این سه الگوریتم نه به‌صورت تصادفی، بلکه بر اساس یک طراحی مفهومی ساختارمند و با هدف پوشش سه پارادایم اصلی در طبقه‌بندی داده‌های سنجنش از دور انجام شده است: (۱) روش‌های یادگیری ماشین نظارت‌شده، (۲) روش‌های آماری کلاسیک، و (۳) روش‌های بدون نظارت. این رویکرد امکان تحلیل نظام‌مند نقاط قوت و ضعف هر دسته از الگوریتم‌ها را در یک بستر داده‌ای مشترک فراهم کرده و به ارائه نتایجی قابل اتکا و قابل تعمیم در مطالعات LULC کمک می‌کند.

۲-۴-۲-۱. جنگل تصادفی (Random Forest)

الگوریتم جنگل تصادفی یک روش یادگیری جمعی (Ensemble Learning) مبتنی بر مجموعه‌ای از درخت‌های تصمیم‌گیری است که هر یک از آن‌ها با استفاده از نمونه‌برداری تصادفی با جایگزینی



شکل ۲. روند کلی انجام کار

۲-۴-۲. الگوریتم‌های طبقه‌بندی

در این پژوهش، انتخاب الگوریتم‌های طبقه‌بندی به‌صورت هدفمند و بر مبنای یک راهبرد مقایسه‌ای میان پارادایم‌های مختلف یادگیری انجام شده است تا امکان ارزیابی جامع رفتار مدل‌ها در مواجهه با داده‌های چندمنبعی و محیط‌های شهری پیچیده فراهم گردد. برخلاف بسیاری از مطالعات که صرفاً بر یک الگوریتم خاص تمرکز دارند یا مقایسه‌ای محدود میان روش‌های هم‌خانواده ارائه می‌دهند، در این تحقیق تلاش شده است نمایندگانی از سه رویکرد متمایز شامل یادگیری نظارت‌شده مبتنی بر یادگیری ماشین، روش‌های کلاسیک آماری و الگوریتم‌های بدون نظارت در یک چارچوب یکسان و با شرایط داده‌ای مشابه مورد بررسی قرار گیرند.

در این راستا، الگوریتم جنگل تصادفی (Random Forest) به‌عنوان یک روش یادگیری جمعی مبتنی بر درخت‌های تصمیم انتخاب شده است. این الگوریتم به دلیل توانایی بالا در مدل‌سازی روابط غیرخطی، تحمل مناسب نسبت به داده‌های نویزی و چندبعدی، و مقاومت در برابر بیش‌برازش، به‌طور گسترده در مطالعات سنجنش از دور به‌عنوان یکی از روش‌های دقیق و پایدار در طبقه‌بندی LULC شناخته می‌شود. همچنین، قابلیت استفاده همزمان از تعداد زیادی

۲-۴-۳. الگوریتم میانگین گیری (K (K-means

الگوریتم K-means یک روش بدون نظارت مبتنی بر خوشه بندی است که داده ها را بر اساس شباهت طیفی به K خوشه تقسیم می کند. این الگوریتم با انتخاب مراکز اولیه خوشه ها آغاز شده و به صورت تکراری، هر پیکسل را به نزدیک ترین مرکز خوشه اختصاص داده و سپس مراکز خوشه ها را به روزرسانی می کند. [۳۴] این فرآیند تا زمانی ادامه می یابد که تغییرات مراکز خوشه ها یا تخصیص پیکسل ها به حداقل برسد. با وجود سادگی و سرعت محاسباتی بالا، این الگوریتم به مقداردهی اولیه مراکز و انتخاب تعداد خوشه ها (K) حساس بوده و ممکن است به نتایج متفاوتی همگرا شود.

در این پژوهش، هدف از به کارگیری الگوریتم K-means صرفاً انجام یک طبقه بندی بدون نظارت نیست، بلکه این روش به عنوان یک ابزار تحلیلی برای بررسی ساختار ذاتی داده ها و ارزیابی میزان تفکیک پذیری طبیعی کلاس های کاربری زمین در فضای ویژگی های استخراج شده به کار گرفته شده است. در واقع، K-means در این مطالعه نقش یک مرجع بدون نظارت (Unsupervised Reference) را ایفا می کند تا مشخص شود تا چه حد داده های طیفی و شاخص های مشتق شده از تصاویر Sentinel-2 به تنهایی قادر به ایجاد جدایش نسبی بین کلاس ها هستند، بدون آنکه از داده های آموزشی استفاده شود. بنابراین، هدف از مقایسه این الگوریتم با روش های نظارت شده، سنجش عملکرد رقابتی نیست، بلکه تحلیل تفاوت میان رویکردهای مبتنی بر داده های برچسب دار و ساختارهای ذاتی داده ها در یک چارچوب یکسان است.

$$J = \sum_{j=1}^K \sum_{x_i \in C_j} \|x_i - \mu_j\|^2 \quad (3)$$

۲-۴-۳. ارزیابی دقت

در مطالعات طبقه بندی داده های سنجش از دور، ارزیابی صحت نتایج از اهمیت بالایی برخوردار است. صحت طبقه بندی معمولاً با استفاده از معیارهای کمی محاسبه می شود. این معیارها از ماتریس درهم ریختگی (Confusion Matrix) استخراج می شوند که در آن

(Bootstrap Sampling) از داده های آموزشی ساخته می شوند [۳۶]. در این روش، هر درخت به طور مستقل آموزش داده شده و در نهایت، تصمیم نهایی از طریق رأی گیری اکثریت میان درخت ها تعیین می شود.

این ساختار موجب کاهش واریانس مدل، افزایش پایداری و بهبود قدرت تعمیم پذیری در داده های پیچیده می شود. همچنین استفاده از نمونه های خارج از کیسه ۱۰ (Out-of-Bag) امکان ارزیابی داخلی عملکرد مدل را بدون نیاز به مجموعه اعتبارسنجی جداگانه فراهم می سازد. پارامترهای مهم این الگوریتم شامل تعداد درخت ها (Ntrees)، حداکثر عمق درخت ها (Dmax)، حداقل تعداد نمونه در هر گره (Nmin) و تعداد ویژگی های انتخاب شده در هر تقسیم (mtry) هستند که تنظیم آن ها نقش مهمی در دقت نهایی مدل دارد.

$$y^{\wedge} = \text{mode}\{h_1(x), h_2(x), \dots, h_n(x)\} \quad (1)$$

۲-۴-۲. حداکثر احتمال (Maximum Likelihood)

الگوریتم حداکثر احتمال یکی از روش های کلاسیک طبقه بندی نظارت شده است که بر پایه نظریه احتمال و فرض توزیع نرمال چندمتغیره داده ها عمل می کند. در این روش، هر پیکسل به کلاسی اختصاص داده می شود که بیشترین احتمال تعلق به آن کلاس را داشته باشد [۳۵].

این الگوریتم با استفاده از بردار میانگین و ماتریس کوواریانس هر کلاس، احتمال تعلق هر پیکسل را محاسبه می کند. در شرایطی که داده های آموزشی کافی نباشند یا فرض نرمال بودن توزیع داده ها برقرار نباشد، دقت این روش کاهش یافته و حساسیت آن نسبت به نویز افزایش می یابد.

(۲)

$$P(\omega_i | x) = \frac{P(x | \omega_i)P(\omega_i)}{P(x)}$$

¹⁰ Out-of-Bag

۴-۳-۴-۲. ضریب کاپا (Kappa Coefficient)

این معیار برای ارزیابی توافق بین داده‌های طبقه‌بندی شده و داده‌های مرجع به کار می‌رود و شانس تصادفی را نیز در نظر می‌گیرد. مقدار کاپا بین ۱- تا ۱ تغییر می‌کند که مقادیر نزدیک به ۱ نشان‌دهنده توافق بالا و مقادیر نزدیک به صفر نشان‌دهنده توافق ضعیف یا تصادفی است.

$$\kappa = \frac{e_o - p}{e_e p - 1} \quad (7)$$

در این فرمول، e_o دقت مشاهده شده و e_e دقت مورد انتظار ناشی از شانس تصادفی است. ضریب کاپا نسبت به دقت کلی حساس‌تر است، زیرا احتمال توافق تصادفی را اصلاح می‌کند و معیار معتبرتری برای مقایسه مدل‌های مختلف ارائه می‌دهد.

۳. نتایج و ارزیابی

برای مقایسه عملکرد الگوریتم‌های طبقه‌بندی، مجموعه‌ای از معیارهای ارزیابی شامل دقت کلی، ضریب کاپا، دقت تولیدکننده، دقت کاربر و مساحت کلاس‌ها در نظر گرفته شد. در این میان، دقت کلی بیانگر نسبت نمونه‌های به‌درستی طبقه‌بندی شده به کل داده‌ها بوده و ضریب کاپا میزان توافق نتایج طبقه‌بندی با داده‌های مرجع را با لحاظ خطای تصادفی نشان می‌دهد.

داده‌های مورد استفاده در این پژوهش شامل تصاویر ماهواره‌ای Sentinel-2 و داده‌های زمینی مرجع مربوط به مهرماه ۱۴۰۳ است. این داده‌ها به‌عنوان ورودی مشترک برای تمامی الگوریتم‌های مورد بررسی مورد استفاده قرار گرفته‌اند.

برای اجرای الگوریتم جنگل تصادفی، تنظیمات مدل شامل ۱۰۰ درخت تصمیم ($n_estimators = 100$)، حداکثر عمق ۲۰ ($max_depth = 20$)، حداقل تعداد نمونه در هر برگ ۲ ($min_samples_leaf = 2$) و انتخاب ویژگی‌ها بر اساس ریشه دوم تعداد کل ویژگی‌ها $max_features = \sqrt{\text{تعداد کل ویژگی‌ها}}$ در نظر گرفته شده است.

نتایج حاصل از اجرای الگوریتم‌ها در شکل‌های ۳، ۴ و ۵ و همچنین جداول ۲، ۳ و ۴ ارائه شده‌اند که به تفکیک هر روش طبقه‌بندی را نمایش می‌دهند.

نمونه‌های پیش‌بینی شده برای هر کلاس با مقادیر مرجع مقایسه می‌شوند.

۴-۳-۴-۲. دقت کلی (Overall Accuracy)

این شاخص بیانگر درصد پیکسل‌هایی است که به‌درستی به کلاس‌های مربوطه تخصیص داده شده‌اند و معیاری ساده برای ارزیابی کلی عملکرد مدل طبقه‌بندی است.

$$Accuracy = \frac{\sum TP + TN}{\sum TP + FP + FN + TN} \quad (4)$$

در این فرمول، TP تعداد نمونه‌های درست پیش‌بینی شده برای کلاس مثبت، TN تعداد نمونه‌های درست پیش‌بینی شده برای کلاس منفی، FP نمونه‌های پیش‌بینی شده نادرست به عنوان کلاس مثبت و FN نمونه‌های کلاس مثبت که نادرست پیش‌بینی شده‌اند، هستند.

۴-۳-۴-۲. دقت تولیدکننده (Producer's Accuracy)

این معیار نشان می‌دهد که یک کلاس چقدر به‌درستی تشخیص داده شده است.

$$PA = \frac{TP}{TP + FN} \quad (5)$$

در واقع، دقت تولیدکننده بیان می‌کند که چه سهمی از نمونه‌های واقعی یک کلاس خاص، توسط الگوریتم به درستی طبقه‌بندی شده‌اند.

۴-۳-۴-۲. دقت کاربر (User's Accuracy)

نشان می‌دهد که چند درصد از پیکسل‌های پیش‌بینی شده برای یک کلاس خاص واقعاً به آن کلاس تعلق دارند. به عبارت دیگر، دقت کاربر میزان اطمینان یک کاربر را از صحت طبقه‌بندی پیکسل‌ها ارائه می‌دهد.

$$UA = \frac{TP}{TP + FP} \quad (6)$$

بر اساس شکل ۴ و جدول ۳، نتایج حاصل از اجرای الگوریتم Random Forest نشان می‌دهد که این روش به صحت کلی ۹۸٪ و ضریب کاپا ۹۵٪ دست یافته است که بیانگر عملکرد پایدار و سازگار مدل در طبقه‌بندی داده‌ها است.

در بررسی توزیع مکانی کلاس‌ها، نتایج نشان می‌دهد که کلاس خاک با ۶۷۶.۳ هکتار بیشترین سهم را در منطقه مورد مطالعه به خود اختصاص داده است. پس از آن، کلاس‌های ساختمان با ۲۶۹.۷ هکتار، گیاه با ۱۲۰.۸ هکتار و آب با ۱۳.۲ هکتار قرار دارند. این مقادیر در جدول ۳ به تفکیک ارائه شده‌اند و مبنای مقایسه فضایی بین کلاس‌ها در این روش محسوب می‌شوند.

در ارزیابی دقت‌های کلاسی، نتایج نشان می‌دهد که الگوریتم Random Forest در تمامی کلاس‌ها عملکرد بالایی را ثبت کرده است. به‌طوری‌که کلاس ساختمان دارای دقت کاربر و تولیدکننده برابر با ۱۰۰٪ بوده و کلاس خاک نیز با دقت کاربر ۹۷.۷٪ و دقت تولیدکننده ۱۰۰٪ گزارش شده است. همچنین برای کلاس گیاه، دقت کاربر ۹۷.۶٪ و دقت تولیدکننده ۹۰.۵٪ به دست آمده و در نهایت کلاس آب دارای دقت کاربر ۹۰.۸٪ و دقت تولیدکننده ۹۱.۱٪ بوده است.

به‌طور کلی، نتایج این بخش در شکل ۴ و جدول ۳ نشان‌دهنده خروجی کامل طبقه‌بندی در سطح کلاس‌های مختلف بوده و تمامی شاخص‌های ارزیابی به‌صورت همزمان در تحلیل عملکرد مدل لحاظ شده‌اند.

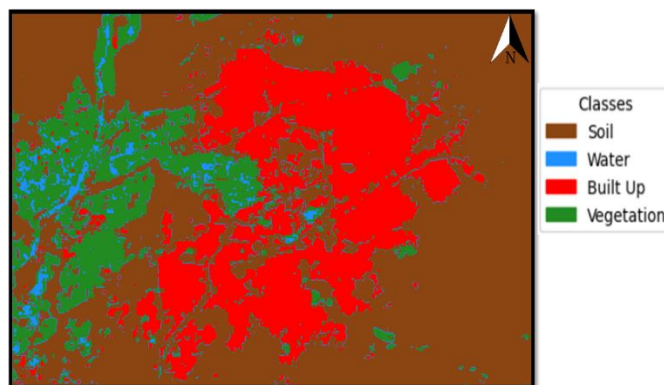


شکل ۴. نقشه حاصل از طبقه‌بندی منطقه نایسر به کمک روش جنگل تصادفی

بر اساس شکل ۳ و جدول ۲، نتایج الگوریتم Maximum Likelihood نشان می‌دهد که صحت کلی برابر با ۹۵٪ و ضریب کاپا حدود ۹۲٪ به دست آمده است. همچنین مساحت محاسبه‌شده برای کلاس‌ها به شرح زیر است: کلاس خاک با ۶۴۶.۷ هکتار، کلاس ساختمان با ۲۷۳.۸ هکتار، کلاس گیاه با ۱۴۲ هکتار و کلاس آب با ۱۸ هکتار.

در بخش دقت‌های کلاسی این الگوریتم، نتایج به‌صورت زیر گزارش شده است:

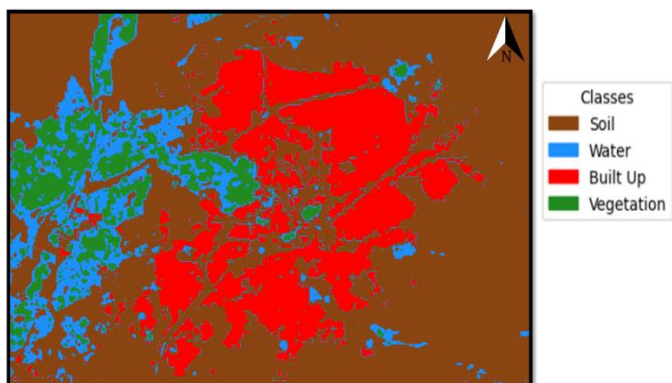
کلاس ساختمان دارای دقت کاربر ۱۰۰٪ و دقت تولیدکننده ۱۰۰٪، کلاس خاک دارای دقت کاربر ۱۰۰٪ و دقت تولیدکننده ۱۰۰٪، کلاس گیاه دارای دقت کاربر ۹۳٪ و دقت تولیدکننده ۸۳٪ و کلاس آب دارای دقت کاربر ۶۶٪ و دقت تولیدکننده ۸۵٪ است.



شکل ۳. نقشه حاصل از طبقه‌بندی منطقه نایسر به کمک روش حداکثر احتمال

جدول ۲. ارزیابی دقت و مساحت کلاس‌ها در روش طبقه‌بندی حداکثر احتمال

کلاس	مساحت (هکتار)	دقت کاربر (%)	دقت تولیدکننده (%)	ضریب کاپا (%)	صحت کلی (%)
ساختمان	۲۷۳.۸	۱۰۰	۱۰۰	۹۲.۷	۹۵
گیاه	۱۴۲	۹۳.۸	۸۳.۵		
آب	۱۸	۶۶.۶۶	۸۵.۷		
خاک	۶۴۶.۶	۱۰۰	۱۰۰		



شکل ۵. نقشه حاصل از طبقه‌بندی منطقه نایسر به کمک روش الگوریتم میانگین‌گیری K (K-means)

جدول ۴. ارزیابی دقت و مساحت کلاس‌ها در روش طبقه‌بندی الگوریتم میانگین‌گیری K (K-means)

کلاس	مساحت (هکتار)	دقت کاربر (%)	دقت تولیدکننده (%)	ضریب کاپا (%)	صحت کلی (%)
ساختمان	۲۵۱.۵	۱۰۰	۹۹.۵	۸۷	۹۱
گیاه	۷۷.۹	۷۳.۲	۸۸.۵		
آب	۱۱۰.۷	۵۵.۵	۶۰.۵		
خاک	۶۴۰	۱۰۰	۱۰۰		

۴. بحث و نتیجه‌گیری

در این پژوهش، عملکرد سه الگوریتم طبقه‌بندی شامل حداکثر احتمال (Maximum Likelihood)، جنگل تصادفی (Random Forest) و K-means در طبقه‌بندی تصاویر Sentinel-2 مورد ارزیابی قرار گرفت. نتایج نشان داد که الگوریتم Random Forest با دستیابی به بالاترین دقت کلی و ضریب کاپا، نسبت به سایر روش‌ها عملکرد برتری داشته و توانسته است تفکیک دقیق‌تری از کلاس‌های کاربری اراضی ارائه دهد.

مقایسه این نتایج با مطالعات پیشین نشان می‌دهد که روند کلی به‌دست‌آمده در این تحقیق با ادبیات موجود در حوزه سنجش از دور هم‌راستا است. در اغلب پژوهش‌ها گزارش شده است که الگوریتم‌های مبتنی بر یادگیری ماشین، به‌ویژه Random Forest، به

جدول ۳. ارزیابی دقت و مساحت کلاس‌ها در روش طبقه‌بندی جنگل تصادفی

کلاس	مساحت (هکتار)	دقت کاربر (%)	دقت تولیدکننده (%)	ضریب کاپا (%)	صحت کلی (%)
ساختمان	۲۶۹.۷	۱۰۰	۱۰۰	۹۵	۹۸
گیاه	۱۲۰.۸	۹۷/۶	۹۰.۵		
آب	۱۳.۲	۹۰/۸	۹۱/۱		
خاک	۶۷۶.۳	۹۷/۷	۱۰۰		

مطابق شکل ۵ و جدول ۴، نتایج حاصل از پیاده‌سازی الگوریتم K-means به‌عنوان یک روش بدون‌ناظر در منطقه مورد مطالعه ارائه شده است. در این روش، صحت کلی طبقه‌بندی برابر با ۹۱ درصد و ضریب کاپا معادل ۰.۸۷ به‌دست آمده است که نشان‌دهنده میزان توافق قابل قبول بین نتایج طبقه‌بندی و داده‌های مرجع است.

بر اساس خروجی طبقه‌بندی، مساحت کلاس‌های کاربری و پوشش زمین به‌صورت زیر برآورد شده است: کلاس «ساختمان» ۲۵۱.۵ هکتار، کلاس «گیاه» ۷۷.۹ هکتار، کلاس «آب» ۱۱۰.۷ هکتار و کلاس «خاک» ۶۴۰ هکتار.

در ارزیابی دقت‌های کلاسی، برای کلاس «ساختمان» دقت کاربر برابر با ۱۰۰ درصد و دقت تولیدکننده ۹۹.۵ درصد، برای کلاس «گیاه» دقت کاربر ۷۳.۲ درصد و دقت تولیدکننده ۸۸.۵ درصد، برای کلاس «آب» دقت کاربر ۵۵.۵ درصد و دقت تولیدکننده ۶۰.۵ درصد و برای کلاس «خاک» دقت کاربر و تولیدکننده هر دو برابر با ۱۰۰ درصد محاسبه شده است.

به‌طور کلی، نتایج نشان می‌دهد که عملکرد الگوریتم در کلاس‌های دارای تمایز طیفی بیشتر (مانند ساختمان و خاک) بالاتر بوده است، در حالی که در کلاس‌های دارای شباهت طیفی بیشتر (مانند آب و گیاه) کاهش نسبی دقت مشاهده می‌شود.

این نتایج در جدول ۴ و نقشه طبقه‌بندی‌شده در شکل ۵ ارائه شده‌اند.

Maximum Likelihood و ضعف K-means، نشان‌دهنده وجود برخی تفاوت‌های قابل توجه در میزان دقت و رفتار کلاس‌های مختلف نیز می‌باشد. این تفاوت‌ها عمدتاً ناشی از ویژگی‌های مکانی منطقه مورد مطالعه، کیفیت داده‌های Sentinel-2، ترکیب شاخص‌های طیفی مانند NDVI، NDBI و MNDWI و همچنین افزودن داده‌های کمکی توپوگرافی است. در نتیجه، این مطالعه تأکید می‌کند که عملکرد الگوریتم‌های طبقه‌بندی به شدت وابسته به شرایط محیطی و ساختار ویژگی‌های ورودی است و نمی‌توان یک الگوریتم را به صورت مطلق در تمامی شرایط برتر دانست.

به طور کلی، نتایج این تحقیق نشان می‌دهد که Random Forest به دلیل انعطاف‌پذیری بالا، توانایی تعمیم مناسب و قابلیت پردازش داده‌های پیچیده چندبعدی، گزینه‌ای بسیار کارآمد برای طبقه‌بندی کاربری اراضی در مناطق شهری محسوب می‌شود. در عین حال، استفاده از داده‌های چندمنبعی و شاخص‌های طیفی می‌تواند موجب بهبود عملکرد سایر الگوریتم‌ها شده و دقت کلی طبقه‌بندی را افزایش دهد. این یافته‌ها بر اهمیت انتخاب هوشمندانه الگوریتم و طراحی مناسب فضای ویژگی در مطالعات سنجش از دور تأکید دارند.

دلیل ساختار مبتنی بر مجموعه‌ای از درخت‌های تصمیم، توانایی بالایی در مدل‌سازی روابط غیرخطی، مدیریت داده‌های نویزی و بهره‌گیری هم‌زمان از ویژگی‌های چندبعدی دارند. این ویژگی‌ها موجب شده است که این الگوریتم در بسیاری از مطالعات نسبت به روش‌های کلاسیک آماری مانند Maximum Likelihood عملکرد بهتری داشته باشد. در مطالعه حاضر نیز همین الگو مشاهده شد، با این تفاوت که میزان اختلاف دقت بین این دو روش تا حدی تحت تأثیر ویژگی‌های طیفی منطقه مطالعه و استفاده از شاخص‌های کمکی بوده است.

در مقابل، الگوریتم Maximum Likelihood اگرچه همچنان به عنوان یکی از روش‌های پایه و پرکاربرد در طبقه‌بندی تصاویر سنجش از دور شناخته می‌شود، اما نتایج این تحقیق نشان داد که عملکرد آن نسبت به Random Forest محدودتر است. این موضوع با نتایج مطالعات پیشین که بر حساسیت این روش به فرض نرمال بودن توزیع داده‌ها و همپوشانی طیفی کلاس‌ها تأکید دارند، هم‌خوانی دارد. با این حال، استفاده از داده‌های Sentinel-2 و شاخص‌های طیفی در این مطالعه موجب شد که این الگوریتم در برخی کلاس‌های با تفکیک‌پذیری مناسب، نتایج قابل قبولی ارائه دهد که نشان‌دهنده قابلیت نسبی آن در شرایط داده‌ای بهینه‌تر است.

در خصوص الگوریتم K-means، نتایج به دست آمده نشان داد که این روش به عنوان یک الگوریتم بدون‌ناظر، نسبت به روش‌های نظارت‌شده دقت پایین‌تری دارد. این یافته با بخش عمده‌ای از ادبیات پژوهش هم‌راستا است که بیان می‌کنند روش‌های خوشه‌بندی بدون نظارت در محیط‌هایی با تنوع طیفی بالا و همپوشانی کلاس‌ها، توانایی محدودی در تفکیک دقیق دارند. با این وجود، K-means توانست الگوی کلی ساختار مکانی کاربری‌ها را شناسایی کند و در کلاس‌هایی با تفکیک‌پذیری طیفی بالا مانند ساختمان و خاک عملکرد نسبتاً پایدارتری نسبت به کلاس‌های پیچیده‌تر مانند آب و پوشش گیاهی نشان دهد. این موضوع نشان می‌دهد که علی‌رغم محدودیت‌ها، این الگوریتم می‌تواند در تحلیل‌های اولیه و نبود داده‌های آموزشی نقش مکمل ایفا کند.

مقایسه نتایج این مطالعه با تحقیقات پیشین علاوه بر تأیید روند کلی عملکرد الگوریتم‌ها برتری Random Forest، عملکرد متوسط

- sensing: An applied review. *International journal of remote sensing* 2018, 39(9):2784-2817.
- [10] Gorelick N, Hancher M, Dixon M, Ilyushchenko S, Thau D, Moore R: Google Earth Engine: Planetary-scale geospatial analysis for everyone. *Remote sensing of Environment* 2017, 202:18-27.
- [11] Kempeneers P, Kliment T, Marletta L, Soille P: Parallel Processing Strategies for Geospatial Data in a Cloud Computing Infrastructure. *Remote Sensing* 2022, 14(2):398.
- [12] Liu L, Xiao X, Qin Y, Wang J, Xu X, Hu Y, Qiao Z: Mapping cropping intensity in China using time series Landsat and Sentinel-2 images and Google Earth Engine. *Remote Sensing of Environment* 2020, 239:111624.
- [13] Hermosilla T, Wulder MA, White JC, Coops NC, Hobart GW: Disturbance-informed annual land cover classification maps of Canada's forested ecosystems for a 29-year Landsat time series. *Canadian Journal of Remote Sensing* 2018, 44(1):67-87.
- [14] Pflugmacher D, Rabe A, Peters M, Hostert P: Mapping pan-European land cover using Landsat spectral-temporal metrics and the European LUCAS survey. *Remote sensing of environment* 2019, 221:583-595.
- [15] McCarty DA, Kim HW, Lee HK: Evaluation of light gradient boosted machine learning technique in large scale land use and land cover classification. *Environments* 2020, 7(10):84.
- [17] Thonfeld F, Steinbach S, Muro J, Kirimi F: Long-term land use/land cover change assessment of the Kilombero catchment in Tanzania using random forest classification and robust change vector analysis. *Remote sensing* 2020, 12(7):1057.
- [18] Santos LA, Ferreira K, Picoli M, Camara G, Zurita-Milla R, Augustijn E-W: Identifying spatiotemporal patterns in land use and cover samples from satellite image time series. *Remote Sensing* 2021, 13(5):974.
- [19] Woznicki SA, Baynes J, Panlasigui S, Mehaffey M, Neale A: Development of a spatially complete floodplain map of the conterminous United States using random forest. *Science of the total environment* 2019, 647:942-953.
- [20] Betts MG, Wolf C, Ripple WJ, Phalan B, Millers KA, Duarte A, Butchart SH, Levi T: Global forest

References

- [1] Esfandeh S, Danehkar A, Salmanmahiny A, Sadeghi SMM, Marcu MV: Climate change risk of urban growth and land use/land cover conversion: An in-depth review of the recent research in Iran. *Sustainability* 2021, 14(1):338.
- [2] Yao Y, Yan X, Luo P, Liang Y, Ren S, Hu Y, Han J, Guan Q: Classifying land-use patterns by integrating time-series electricity data and high-spatial resolution remote sensing imagery. *International Journal of Applied Earth Observation and Geoinformation* 2022, 106:102664.
- [3] Schulz D, Yin H, Tischbein B, Verleysdonk S, Adamou R, Kumar N: Land use mapping using Sentinel-1 and Sentinel-2 time series in a heterogeneous landscape in Niger, Sahel. *ISPRS Journal of Photogrammetry and Remote Sensing* 2021, 178:97-111.
- [4] Steinhausen MJ, Wagner PD, Narasimhan B, Waske B: Combining Sentinel-1 and Sentinel-2 data for improved land use and land cover mapping of monsoon regions. *International journal of applied earth observation and geoinformation* 2018, 73:595-604.
- [5] Ienco D, Interdonato R, Gaetano R, Minh DHT: Combining Sentinel-1 and Sentinel-2 Satellite Image Time Series for land cover mapping via a multi-source deep learning architecture. *ISPRS Journal of Photogrammetry and Remote Sensing* 2019, 158:11-22.
- [6] Clerici N, Valbuena Calderón CA, Posada JM: Fusion of Sentinel-1A and Sentinel-2A data for land cover mapping: a case study in the lower Magdalena region, Colombia. *Journal of Maps* 2017, 13(2):718-726.
- [7] Wulder MA, White JC, Loveland TR, Woodcock CE, Belward AS, Cohen WB, Fosnight EA, Shaw J, Masek JG, Roy DP: The global Landsat archive: Status, consolidation, and direction. *Remote Sensing of Environment* 2016, 185:271-283.
- [8] Guo J, Huang C, Hou J: A scalable computing resources system for remote sensing big data processing using geopyspark based on spark on k8s. *Remote Sensing* 2022, 14(3):521.
- [9] Maxwell AE, Warner TA, Fang F: Implementation of machine-learning classification in remote

- environmental change. *Trends in ecology & evolution* 2005, 20(9):503-510.
- [30] Szabo S, Gácsi Z, Balazs B: Specific features of NDVI, NDWI and MNDWI as reflected in land cover categories. 2016.
- [31] Huete AR: A soil-adjusted vegetation index (SAVI). *Remote sensing of environment* 1988, 25(3):295-309.
- [32] Zhang Y, Odeh IO, Han C: Bi-temporal characterization of land surface temperature in relation to impervious surface area, NDVI and NDBI, using a sub-pixel image analysis. *International Journal of Applied Earth Observation and Geoinformation* 2009, 11(4):256-264.
- [32] Hassan-Esfahani L, Torres-Rua A, Jensen A, McKee M: Assessment of surface soil moisture using high-resolution multi-spectral imagery and artificial neural networks. *Remote Sensing* 2015, 7(3):2627-2646.
- [33] Friedman J: The elements of statistical learning: Data mining, inference, and prediction. (No Title) 2009.
- [34] Belgiu M, Drăguț L: Random forest in remote sensing: A review of applications and future directions. *ISPRS Journal of Photogrammetry and Remote Sensing* 2016, 114:24-31.
- [35] Maxwell AE, Warner TA, Fang F: Implementation of machine learning classification in remote sensing: An applied review. *International Journal of Remote Sensing* 2018, 39(9):2784-2817.
- [36] Pelletier C, Valero S, Inglada J, Champion N, Dedieu G: Assessing the robustness of random forests to map land cover with high resolution satellite image time series over large areas. *Remote Sensing of Environment* 2016, 187:156-168.
- [37] Pham TD, Yokoya N, Xia J, Ha NT, Xia GS: Comparison of machine learning methods for land use/land cover classification using multi-source remote sensing data. *Remote Sensing* 2020, 12(10):1683.
- [38] Li W, Fu H, Yu L, Cracknell A: Deep learning based oil palm tree detection and counting for high-resolution remote sensing images. *Remote Sensing* 2017, 9(1):22.
- [39] Khatami R, Mountrakis G, Stehman SV: A meta-analysis of remote sensing research on supervised pixel-based land-cover image classification processes: General guidelines for practitioners and loss disproportionately erodes biodiversity in intact landscapes. *Nature* 2017, 547(7664):441-444.
- [21] Smits P, Dellepiane S, Schowengerdt R: Quality assessment of image classification algorithms for land-cover mapping: A review and a proposal for a cost-based approach. *International journal of remote sensing* 1999, 20(8):1461-1486.
- [22] optimisation of image acquisition for land cover classification with Random Forest and MODIS time-series. *International Journal of Applied Earth Observation and Geoinformation* 2015, 34:136-146.
- [23] Rodriguez-Galiano VF, Ghimire B, Rogan J, Chica-Olmo M, Rigol-Sanchez JP: An assessment of the effectiveness of a random forest classifier for land-cover classification. *ISPRS journal of photogrammetry and remote sensing* 2012, 67:93-104.
- [24] Shao Y, Lunetta RS: Comparison of support vector machine, neural network, and CART algorithms for the land-cover classification using limited training data points. *ISPRS Journal of Photogrammetry and Remote Sensing* 2012, 70:78-87.
- [25] Goodin DG, Anibas KL, Bezymennyi M: Mapping land cover and land use from object-based classification: An example from a complex agricultural landscape. *International Journal of Remote Sensing* 2015, 36(18):4702-4723.
- [26] Mansaray LR, Wang F, Huang J, Yang L, Kanu AS: Accuracies of support vector machine and random forest in rice mapping with Sentinel-1A, Landsat-8 and Sentinel-2A datasets. *Geocarto International* 2020, 35(10):1088-1108.
- [27] Hirayama H, Sharma RC, Tomita M, Hara K: Evaluating multiple classifier system for the reduction of salt-and-pepper noise in the classification of very-high-resolution satellite images. *International journal of remote sensing* 2019, 40(7):2542-2557.
- [28] Colditz RR: An evaluation of different training sample allocation schemes for discrete and continuous land cover classification using decision tree-based algorithms. *Remote Sensing* 2015, 7(8):9655-9681.
- [29] Pettorelli N, Vik JO, Mysterud A, Gaillard J-M, Tucker CJ, Stenseth NC: Using the satellite-derived NDVI to assess ecological responses to

future research. Remote Sensing of Environment
2016, 177:89-100.