

Performance analysis of Support Vector Machine, Random Forest, and Maximum Likelihood algorithms in land use classification of the metropolitan area of Mashhad

Sajedeh Baghban, Mohammad Rahim Rahnema*, Mohammad Ajza Shokuhi, Hossein Vahidi
Dep of Geography, Faculty of Literature and Human Sciences, Ferdowsi University of Mashhad,
Mashhad, Iran

* Corresponding Author: rahnema@um.ac.ir

Abstract

Introduction: Considering that the value and usability of any map produced from satellite images depend on its accuracy, evaluating the accuracy of satellite image classification methods is of great importance. Therefore, this research aims to analyse the performance of Support Vector Machine (SVM), Random Forest (RF), and Maximum Likelihood Classification (MLC) algorithms in identifying land use and land cover (LULC) in the metropolitan area of Mashhad. Numerous algorithms have been developed for satellite image classification to date, and their performance varies under different conditions. For this reason, this study first identifies the most commonly used algorithms through a review of previous research, and then, by assessing the characteristics of various classifiers, selects the three algorithms: Support Vector Machine, Random Forest, and Maximum Likelihood. There are various studies regarding the performance of different classification algorithms, each yielding different results. Given that multiple studies have shown that LULC mapping accuracy is related to time and location, and that each of these studies has emphasized the accuracy of different algorithms, their results cannot be generalized to the geographical conditions of Iran. On the other hand, there has not been sufficient research in the geomorphological conditions of Iran to assess the accuracy of classification algorithms, and most studies validating these algorithms have been conducted in case studies outside of Iran. Therefore, considering the differences in algorithm results under various conditions, examining the accuracy and performance of these algorithms focusing on the extensive and diverse metropolitan area of Mashhad may yield novel and noteworthy findings.

Materials and Methods: The present research is applied in terms of purpose and descriptive-analytical in terms of nature. Data collection in this study has been conducted through a documentary-library method. In this study, images from the OLI sensor on the Landsat 8 satellite were used. The classification of satellite images was performed in two stages: preprocessing and processing. After assessing the accuracy of the classification using the Kappa coefficient, confusion matrix, coefficient of variation, and User's accuracy and Producer's accuracy coefficients, the best algorithm for classifying land uses in the metropolitan area of Mashhad was determined in five classes: 1- Built-up areas, 2- Barren land, 3- Mountainous areas, 4- Green spaces, and 5- Water bodies. **Results and Discussion:** The results from the evaluation of standard deviation (SD) and coefficient of variation (CV) regarding the area share percentage in a LULC class by various algorithms indicate that barren lands were classified with higher accuracy, while water bodies and green spaces were classified with lower accuracy. The examination of U_Accuracy and P_Accuracy coefficients shows that the overall accuracy of the classification for all studied algorithms falls within the range of good to excellent. However, a more detailed examination of these algorithms reveals that the greatest challenge in class identification lies in built-up areas, mountainous regions, and green spaces, whereas the identification of barren lands faces fewer challenges. The Kappa coefficient and analyses based on the confusion matrix also demonstrate the variation in accuracy among each LULC

classifier. The differences in the accuracy of the classifiers used are marginal, but these slight variations hold significant importance in the context of LULC planning. Given that these marginal differences are evident in sensitive land uses such as built-up areas and green spaces, selecting an algorithm with the highest accuracy and lowest error is of special importance.

Conclusion: The results of the Kappa coefficient evaluation and confusion matrix analyses indicate that the SVM approach has greater overall accuracy and a higher Kappa coefficient compared to RF and MLC methods. Specifically, the algorithms achieved overall accuracies of ۰,۹۳, ۰,۸۸, and ۰,۸۰, respectively. Therefore, Support Vector Machine demonstrates the highest accuracy and least error among the studied classifiers. Considering that numerous studies have shown that LULC mapping accuracy is related to time and location, it is suggested that future research analyse the accuracy of classifiers under different morphoclimatic and geomorphic conditions.

تحلیل کارایی الگوریتم‌های ماشین بردار پشتیبان، جنگل تصادفی و حداکثر احتمال در شناسایی کاربری اراضی منطقه کلان‌شهری مشهد

ساجده باغبان، محمد رحیم رهنما*، محمد اجزاء شکوهی، حسین وحیدی

گروه جغرافیا، دانشکده ادبیات و علوم انسانی، دانشگاه فردوسی مشهد، مشهد، ایران

*نویسنده عهده دار مکاتبات: rahnama@um.ac.ir

چکیده

سابقه و هدف: با توجه به این که ارزش و قابلیت استفاده از هر نقشه تولید شده از تصاویر ماهواره‌ای به درجه صحت آن بستگی دارد، ارزیابی صحت روش طبقه‌بندی تصاویر ماهواره‌ای از اهمیت بالایی برخوردار است. لذا این پژوهش با هدف تحلیل کارایی الگوریتم‌های ماشین بردار پشتیبان (SVM)، جنگل تصادفی (RF) و حداکثر احتمال (MLC) در شناسایی کاربری و پوشش اراضی (LULC) منطقه کلان‌شهری مشهد انجام شده است. الگوریتم‌های بسیار زیادی به منظور طبقه‌بندی تصاویر ماهواره‌ای تا به امروز توسعه یافته‌اند که عملکرد آن‌ها در شرایط مختلف، متفاوت است. به همین دلیل، در این پژوهش ابتدا با مروری بر پژوهش‌های پیشین، پرکاربردترین الگوریتم‌ها مورد شناسایی قرار گرفته و سپس با سنجش ویژگی‌های انواع طبقه‌بندی کننده‌ها، سه الگوریتم ماشین بردار پشتیبان، جنگل تصادفی و حداکثر احتمال انتخاب شده است. با توجه به این که مطالعات متعدد نشان داده است که دقت نقشه برداری LULC با زمان و مکان در ارتباط است و هر یک از پژوهش‌های انجام شده نیز بر دقت الگوریتم‌های متفاوتی تاکید کرده‌اند. لذا نتایج آن‌ها برای شرایط جغرافیایی ایران قابل تعمیم نیست. از طرفی در شرایط ژئومورفولوژیک ایران، پژوهش‌های کافی به منظور سنجش دقت الگوریتم‌های طبقه‌بندی انجام نشده است و اغلب مطالعات صحت سنجی الگوریتم‌ها در نمونه‌های موردی خارج از ایران انجام شده است. لذا با توجه به تفاوت نتایج الگوریتم‌ها در شرایط متفاوت، بررسی دقت و عملکرد الگوریتم‌ها با تمرکز بر منطقه وسیع و متنوع کلان‌شهری مشهد، می‌تواند نتایج بدیع و جالب توجهی را به همراه داشته باشد.

مواد و روش‌ها: روش تحقیق حاضر از منظر هدف، کاربردی و از منظر ماهیت، توصیفی-تحلیلی است. گردآوری اطلاعات در این پژوهش به روش اسنادی-کتابخانه‌ای انجام شده است. در این مطالعه تصویر سنجده OLI در ماهواره لندست ۸ تهیه شده است. طبقه‌بندی تصاویر ماهواره‌ای در دو مرحله پیش پردازش و پردازش تصاویر انجام شده و پس از ارزیابی صحت طبقه‌بندی تصاویر با استفاده از ضریب کاپا، ماتریس اختلاط، ضریب تغییرات و ضرایب User's accuracy و Producer's accuracy،

بهترین الگوریتم در طبقه‌بندی کاربری‌های منطقه کلان‌شهری مشهد در ۵ طبقه ۱- مناطق ساخته شده^۱ ۲- اراضی بایر^۲ ۳- مناطق کوهستانی^۳ ۴- فضاهای سبز^۴ و ۵- پهنه‌های آبی^۵ مشخص شد.

نتایج و بحث: نتایج حاصل ارزیابی انحراف معیار (SD) و ضریب تغییرات (CV) درصد سهم مساحت در یک کلاس LULC توسط الگوریتم‌های مختلف نشان می‌دهد که اراضی بایر با دقت بیش‌تر و پهنه‌های آبی و فضاهای سبز با دقت کم‌تری طبقه‌بندی شده‌اند. نتایج بررسی ضرایب $U_Accuracy$ و $P_Accuracy$ نشان می‌دهد که به طور کلی صحت طبقه‌بندی طبقات در تمام الگوریتم‌های مورد مطالعه در بازه بین خوب تا عالی قرار می‌گیرد. اما بررسی دقیق‌تر این الگوریتم‌ها نشان می‌دهد که بیش‌ترین چالش شناسایی طبقه برای مناطق ساخته شده، مناطق کوهستانی و فضاهای سبز است و شناسایی اراضی بایر با چالش کم‌تری مواجه است. ضریب کاپا و تحلیل‌های مبتنی بر ماتریس اختلاط نیز تنوع در دقت هر طبقه‌بندی‌کننده LULC را نشان می‌دهد. تفاوت در دقت طبقه‌بندی‌کننده‌های مورد استفاده جزئی است، اما این تغییرات جزئی اهمیت بسیار مهمی در زمینه برنامه ریزی LULC دارد. با توجه به این که این اختلافات جزئی در کاربری‌های حساسی مانند مناطق ساخته شده و فضاهای سبز دیده می‌شود، لذا انتخاب الگوریتمی با بیش‌ترین دقت و کم‌ترین خطا از اهمیت ویژه‌ای برخوردار است.

نتیجه‌گیری: نتایج بررسی ضریب کاپا و تحلیل‌های مبتنی بر ماتریس اختلاط نشان می‌دهد که رویکرد SVM دارای دقت کلی بیش‌تری و ضریب کاپای بالاتری نسبت به روش‌های RF و MLC است. به‌طوری‌که الگوریتم‌های SVM، RF و MLC به ترتیب دقت کلی معادل ۰/۹۳، ۰/۸۸ و ۰/۸۰ درصد را به دست آورده‌اند. لذا، ماشین بردار پشتیبان بالاترین دقت و کم‌ترین خطا را در بین طبقه‌بندی‌کننده‌های مورد مطالعه دارد. با توجه به این که مطالعات متعدد نشان داد که دقت نقشه برداری LULC با زمان و مکان در ارتباط است. بنابراین، برای تحقیقات آینده، آنالیز دقت طبقه‌بندی‌کننده‌ها برای شرایط مورفوکلیماتیک و ژئومورفیک متفاوت پیشنهاد می‌شود.

واژه‌های کلیدی: سنجش‌ازدور، طبقه‌بندی کاربری اراضی، ماشین بردار پشتیبان، جنگل تصادفی، بیش‌ترین احتمال، مشهد

۱-مقدمه

کاربری اراضی به‌عنوان یکی از اجزاء سیستم شهری، فرایندهای دینامیک مکانی-زمانی را شامل می‌شود که باتوجه‌به ارزش زمین تحت کنترل درآوردن آن از اهمیت ویژه‌ای برخوردار است؛ بنابراین آگاهی از تغییرات و تحولات کاربری اراضی و عوامل مؤثر بر آن در طول یک دوره زمانی برای برنامه‌ریزان و مدیران بسیار حائز اهمیت است و ممکن است هدف مستقیم برنامه‌ریزی فضایی در نظر گرفته شود. به همین دلیل استفاده از روش‌های آشکارسازی تغییرات برای مشخص کردن روند تغییرات باگذشت زمان ضروری به نظر می‌رسد (van Eck, ۲۰۱۶; Aburas et al., ۲۰۱۶; Hossenin et al., ۲۰۱۸; Mohammadi & Akbari, ۲۰۰۸ & koomen, ۲۰۰۸)؛ لذا تهیه لایه پوشش و کاربری اراضی در گذشته و شرایط فعلی و روند تغییرات آن می‌تواند شناخت دقیق از چندوچون تغییرات منطقه ارائه دهد. در این راستا تکنولوژی‌های نوین مانند سنجش‌ازدور و سیستم اطلاعات جغرافیایی ابزارهای مناسبی برای اندازه‌گیری وسعت و الگوی تغییرات در طی زمان فراهم کرده است. نظارت بر رشد و توسعه شهری به یک کاربرد مهم از داده‌های سنجش‌ازدور و سیستم اطلاعات مکانی تبدیل شده است. به‌طوری‌که روند تعیین تغییرات در کاربری

^۱ Built up areas

^۲ Barren lands

^۳ Rocks

^۴ Green spaces

^۵ Water bodies

زمین بر اساس چند دوره داده‌های سنجش‌ازدور امکان‌پذیر است. استفاده از تکنیک‌های سنجش‌ازدور می‌تواند سطح و روند تغییرات پوشش اراضی را با هزینه کمتر و سرعت بیشتر ممکن سازد. تصاویر ماهواره‌ای به دلیل قدرت تفکیک مکانی بالا، تکرارپذیر بودن و قابلیت پردازش‌های کامپیوتری، از ابزارهای مهم در تهیه نقشه‌های کاربری اراضی و نیز مطالعه تغییرات زمانی در مناطق مختلف است که در این روش می‌توان کاربری اراضی را برحسب تصاویر چندزمانه تفکیک کرد و از روش مقایسه بعد از طبقه‌بندی برای پایش تغییرات استفاده کرد و به شناخت مناسبی از چگونگی تغییرات کاربری اراضی دست‌یافت (Esmaili, ۲۰۱۱; Mendoza et al., ۲۰۱۷; De Espindola et al., ۲۰۱۹; & Ashjaei).

اما باتوجه‌به این که ارزش و قابلیت استفاده از هر نقشه تولید شده از تصاویر ماهواره‌ای به درجه صحت آن بستگی دارد، ارزیابی صحت روش طبقه‌بندی تصاویر ماهواره‌ای از اهمیت بالایی برخوردار است. صحت نقشه‌ها وابسته به فاکتورهایی نظیر صحت داده‌های اولیه و روش طبقه‌بندی تصویر ماهواره‌ای است. اصولاً، طبقه‌بندی فرایند دسته‌بندی پیکسل‌ها در تعدادی مشخص از دسته‌بندی بر اساس ارزش‌های بازتابشی ثبت شده در فضای هر تصویر ماهواره‌ای است. در عمل طبقه‌بندی هر کدام از درجه روشنایی‌ها به کالس‌های پوشش اراضی، زمین‌شناسی، کاربری اراضی و دیگر عوارض سطح زمین منتسب می‌شود (Abburu & Golla, ۲۰۱۵; Sowmya et al., ۲۰۱۹).

الگوریتم‌های بسیار زیادی به‌منظور طبقه‌بندی تصاویر ماهواره‌ای تا به امروز توسعه یافته‌اند که عملکرد آن‌ها در شرایط مختلف، متفاوت است؛ لذا در این پژوهش ابتدا با مروری بر پژوهش‌های پیشین، پرکاربردترین الگوریتم‌ها مورد شناسایی قرار گرفتند. بررسی پژوهش‌های ده سال اخیر نشان داد که الگوریتم‌های ماشین بردار پشتیبان^۱، جنگل تصادفی^۲، درخت تصمیم^۳، حداکثر احتمال^۴، نزدیک‌ترین همسایه^۵ و شبکه عصبی مصنوعی^۶ از محبوب‌ترین تکنیک‌های طبقه‌بندی تصاویر به شمار می‌روند. در راستای مقایسه الگوریتم‌های برتر طبقه‌بندی، ویژگی‌های هر یک از آن‌ها مورد ارزیابی قرار گرفت. به‌طور کلی تنوع در رویکردها، کارایی در مسائل مختلف، سازگاری با داده‌های پیچیده، قابلیت تفسیر و تجربه‌های پیشین از جمله مهم‌ترین عواملی است که در انتخاب الگوریتم‌های جنگل تصادفی، بیشترین احتمال و ماشین بردار پشتیبان در این پژوهش تأثیر داشته‌اند.

پژوهش‌های تالاکدر (۲۰۲۰)، سانتارسیرو (۲۰۲۲) و چندرا و بدی (۲۰۲۱) نشان می‌دهد که روش ماشین بردار پشتیبان از دقت بالایی برخوردار است (Talukdar et al., ۲۰۲۰; Santarsiero et al., ۲۰۲۲; Chandra & Bedi, ۲۰۲۱). همچنین الگوریتم جنگل تصادفی در پژوهش‌های مرادی و کسانی (۱۴۰۳)، شی و همکارانش (۲۰۱۹)؛ تازوین و ویبو (۲۰۲۳) و جیاو و همکارانش (۲۰۲۰) مورد تأیید قرار گرفته است (Moradi & KAsani, ۲۰۲۴; Shi et al., ۲۰۱۹; Marudi et al., ۲۰۲۴; Naswin & Wibowo, ۲۰۲۳; Jiao et al., ۲۰۲۰). بررسی پژوهش‌های انجام شده در زمینه طبقه‌بندی کاربری اراضی نشان می‌دهد که الگوریتم حداکثر احتمال، بیشترین فراوانی را در میان الگوریتم‌های مورد استفاده به‌خصوص در پژوهش‌های داخلی دارد (Rahnama, ۲۰۲۱; MahmoudZadeh et al., ۲۰۱۹; PourKhabbaz et al., ۲۰۱۵; Gholam Ali Fard et al., ۲۰۱۸; MirAkhoyrlou & RahimZadegan, ۲۰۱۴; Azizi Ghalati et al., ۲۰۱۲).

در زمینه بررسی عملکرد الگوریتم‌های مختلف طبقه‌بندی پژوهش‌های مختلفی وجود دارد که نتایج هر یک از آن‌ها متفاوت است. در این میان می‌توان به پژوهش جهانبخشی و اختصاصی (۱۳۹۶) اشاره کرد که در آن الگوریتم‌های جنگل تصادفی، ماشین بردار پشتیبان و بیشترین شباهت در طبقه‌بندی انواع کاربری کشاورزی مورد مقایسه قرار گرفته است. در این پژوهش از تصاویر ماهواره لندست ۸ استفاده گردیده است (JahanBakhshi & Ekhtesasi, ۲۰۱۷). همچنین روش‌های جنگل تصادفی و ماشین بردار پشتیبان به‌منظور طبقه‌بندی تصاویر ماهواره سنتینل در پژوهش قدسی و همکاران (۱۳۹۹) برای طبقه‌بندی کاربری اراضی کشاورزی مورد استفاده قرار گرفته است (Ghodsi et al., ۲۰۲۰). تالاکدر و همکارانش (۲۰۲۰)، شش الگوریتم از الگوریتم‌های زیر مجموعه یادگیری ماشین را بررسی کرده‌اند (Talukdar et al., ۲۰۲۰). الجنابی و همکارانش (۲۰۲۴) نیز مانند همین

^۱ Support vector machine(SVM)

^۲ Random Forest(RF)

^۳ Decision Tree(DT)

^۴ Maximum Likelihood(ML)

^۵ k-nearest neighbors(KNN)

^۶ Artificial Neural Networks ANN

پژوهش، الگوریتم‌های جنگل تصادفی، ماشین بردار پشتیبان و بیشترین احتمال را مورد بررسی قرار داده‌اند (Aljanabi et al., ۲۰۲۴).

یکی از مهم‌ترین معایب پژوهش‌هایی که دقت الگوریتم‌ها را با یکدیگر مقایسه کرده‌اند این است که هیچ معیاری در انتخاب الگوریتم‌های مورد مقایسه وجود ندارد. در واقع باتوجه به تعدد الگوریتم‌های موجود مشخص نیست که پژوهشگر چرا چند الگوریتم خاص را برای مقایسه انتخاب کرده است؛ لذا در این پژوهش، به منظور رفع نواقص موجود در پژوهش‌های پیشین، ابتدا، محبوب‌ترین الگوریتم‌ها مورد بررسی قرار گرفته و سپس سعی شده است با سنجش ویژگی‌های انواع طبقه‌بندی‌کننده‌ها، دلایل انتخاب آن‌ها به طور شفاف بیان شود. همچنین باتوجه به این که مطالعات متعدد نشان داده است که دقت نقشه‌برداری LULC با زمان و مکان در ارتباط است و هر یک از پژوهش‌های انجام شده نیز بر دقت الگوریتم‌های متفاوتی تأکید کرده‌اند؛ لذا نتایج آن‌ها برای شرایط جغرافیایی ایران قابل‌تعمیم نیست. از طرفی در شرایط ژئو مورفولوژیک ایران، پژوهش‌های کافی به منظور سنجش دقت الگوریتم‌های طبقه‌بندی انجام نشده است و در بسیاری از مطالعات، صحت‌سنجی الگوریتم‌ها بر روی نمونه‌های موردی خارج از ایران انجام شده است. این امر نشان‌دهنده تفاوت‌های فرهنگی، جغرافیایی و زیست‌محیطی است که می‌تواند بر نحوه عملکرد الگوریتم‌ها تأثیرگذار باشد. لذا، با توجه به این تفاوت‌ها، بررسی دقت و عملکرد الگوریتم‌های مختلف در شرایط خاص و بی‌نظیر کلان‌شهری مشهد، می‌تواند منجر به نتایج جدید و بدیع در حوزه طبقه‌بندی کاربری اراضی شود.

علاوه بر این، در پژوهش‌های قبلی، به چالش‌های خاص مرتبط با طبقه‌بندی کاربری‌های مختلف به‌طور جامع پرداخته نشده است. این چالش‌ها می‌تواند شامل تنوع کاربری‌ها، پیچیدگی الگوهای جغرافیایی و تغییرات مستمر در استفاده از اراضی باشد. با شناسایی و تحلیل این چالش‌ها در الگوریتم‌های مختلف، این تحقیق سعی دارد به پژوهشگران و تصمیم‌گیرندگان کمک نماید تا در مطالعات آتی خود، بر اساس کلاس‌های طبقه‌بندی و ویژگی‌های منطقه‌ای، الگوریتم‌های مناسب‌تری را انتخاب کنند.

همچنین، انتخاب کلان‌شهر مشهد به‌عنوان منطقه مورد مطالعه به‌دلیل تحولات گسترده‌ای که در چند دهه اخیر در آن رخ داده، از اهمیت ویژه‌ای برخوردار است. مشهد به‌عنوان یکی از مهم‌ترین قطب‌های توسعه در شرق کشور، شاهد تغییرات اجتماعی، اقتصادی و زیست‌محیطی قابل توجهی بوده که می‌تواند بر طبقه‌بندی کاربری اراضی تأثیرگذار باشد.

تنوع جغرافیایی و پویایی‌های موجود در این منطقه، از عوامل اصلی است که می‌تواند طبقه‌بندی کاربری‌های مختلف را به چالش بکشد. به‌طوری که اعمال موفق یک روش در چنین منطقه‌ای، می‌تواند قابلیت استفاده از آن را در دیگر محدوده‌های مشابه تضمین کند. در نهایت، با در نظر گرفتن تنوع الگوریتم‌های موجود برای طبقه‌بندی تصاویر ماهواره‌ای و نیاز به بررسی کارایی آن‌ها در شرایط جغرافیایی خاص ایران، این تحقیق با هدف تحلیل کارایی سه الگوریتم پرکاربرد، شامل ماشین بردار پشتیبان، جنگل تصادفی و حداکثر احتمال، در شناسایی کاربری‌های اراضی کلان‌شهری مشهد انجام می‌شود. هدف این است که به‌منظور بهبود روش‌های طبقه‌بندی و ارتقاء مدیریت و برنامه‌ریزی اراضی، نتایج کاربردی و قابل اعتمادی ارائه گردد.

۲- مواد و روش‌ها

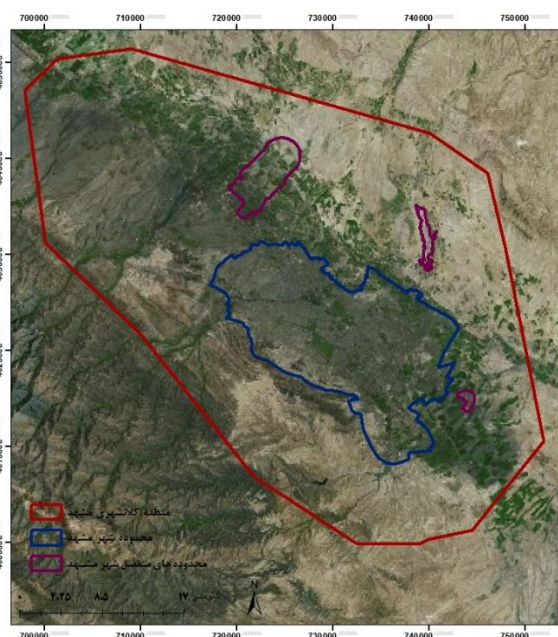
۲-۱- محدوده مورد مطالعه

کلان‌شهر مشهد در $30^{\circ}59'$ تا $35^{\circ}06'$ طول شرقی و $36^{\circ}59'$ تا $35^{\circ}42'$ عرض شمالی واقع شده است و با مساحتی در حدود ۱۸ هزار کیلومتر مربع و با بیش از ۳ میلیون نفر جمعیت (براساس سرشماری سال ۱۳۹۵)، دومین کلان‌شهر ایران به شمار می‌آید (Mashhad municipality, ۲۰۱۶). پویایی منطقه کلان‌شهری مشهد عامل مهمی است که در انتخاب این منطقه به عنوان محدوده مورد مطالعه موثر واقع شده است. عوامل طبیعی و انسانی زیادی بر پویایی محدوده تأثیرگذار هستند. کمربند سبز مشهد ۲۷ کیلومتر طول دارد و در جنوب مشهد قرار دارد که بخشی از آن از میان کوه‌های جنوب شهر می‌گذرد. هر ساله قسمت‌هایی از دامنه‌های این کوه‌ها به زیر ساخت و ساز رفته و از بین می‌روند. شهرستان‌های طرقيه و شاندیز با قرار گرفتن در منطقه کلان‌شهری دومین کلان‌شهر ایران، به عنوان یک عامل انسانی بر پویایی غرب و جنوب غرب منطقه تأثیرگذار بوده‌اند. دامنه‌های شمالی رشته کوه بینالود از جنوب این منطقه می‌گذرد و در توسعه شهرستان‌های طرقيه و شاندیز به عنوان مناطق گردشگری نقش بسزایی دارد، از این رو زمین‌های کشاورزی و باغ‌های دره‌های این رشته کوه در معرض خطر قرار دارند. توسعه

نقش گردشگری در شهرستان های طرقله و شاندىز موجب توسعه کاربرى هاى تفرىحى، گردشگرى و تجارى شده و زمین هاى کشاورزى و کوهستانى منطقه را تحت تأثیر قرار داده است. شدىترین تغییرات در کاربرى زمین در شمال شهر مشهد رخ مى دهد. جریان رودخانه کشف و سرشاخه هاى آن در شمال مشهد باعث توسعه زمین هاى کشاورزى در این قسمت شده است. با توجه به وجود کوه هاى هزار مسجد در شمال شهر مشهد و قرار گرفتن شهرستان هاى طرقله و شاندىز در غرب شهر و به دلیل وجود زمین هاى کشاورزى در شمال، تمرکز روستاها در شمال شهر قابل مشاهده است و با گسترش سکونتگاه هاى انسانى به طور غیر رسمى در شمال، شهر نیز در همین راستا توسعه یافته و زمین هاى کشاورزى را تحت تأثیر قرار داده است. پویایى کاربرى اراضى در این قسمت به گونه اى است که از طرفى، سکونتگاه هاى انسانى در دل اراضى باير و کشاورزى توسعه مى یابند و از طرف دیگر، با افزایش خشک سالى و عدم مدیریت صحیح منابع طبیعى، اراضى کشاورزى بسیارى از بین رفته و تبدیل به اراضى باير مى گردند.

تنوع جغرافیایى دیگر عاملى است که بر انتخاب این منطقه به عنوان محدوده مورد مطالعه نقش داشته است. فضاهاى سبز محدوده از تنوع بالایى برخوردار هستند. به گونه اى که قسمت هاى مختلف فضاهاى سبز با پوشش هاى متفاوتى از جمله کشاورزى، جنگلى، آیش، متراکم و کم تراکم مواجه است. همچنین اراضى باير نیز از جنس هاى متفاوتى از جمله خاک و ماسه هستند. پهنه هاى آبى به صورت پراکنده در محدوده مورد مطالعه قابل مشاهده هستند. علاوه بر این، مناطق صخره اى در جنوب منطقه با دره هاى سبز و اراضى باير آمیخته شده اند.

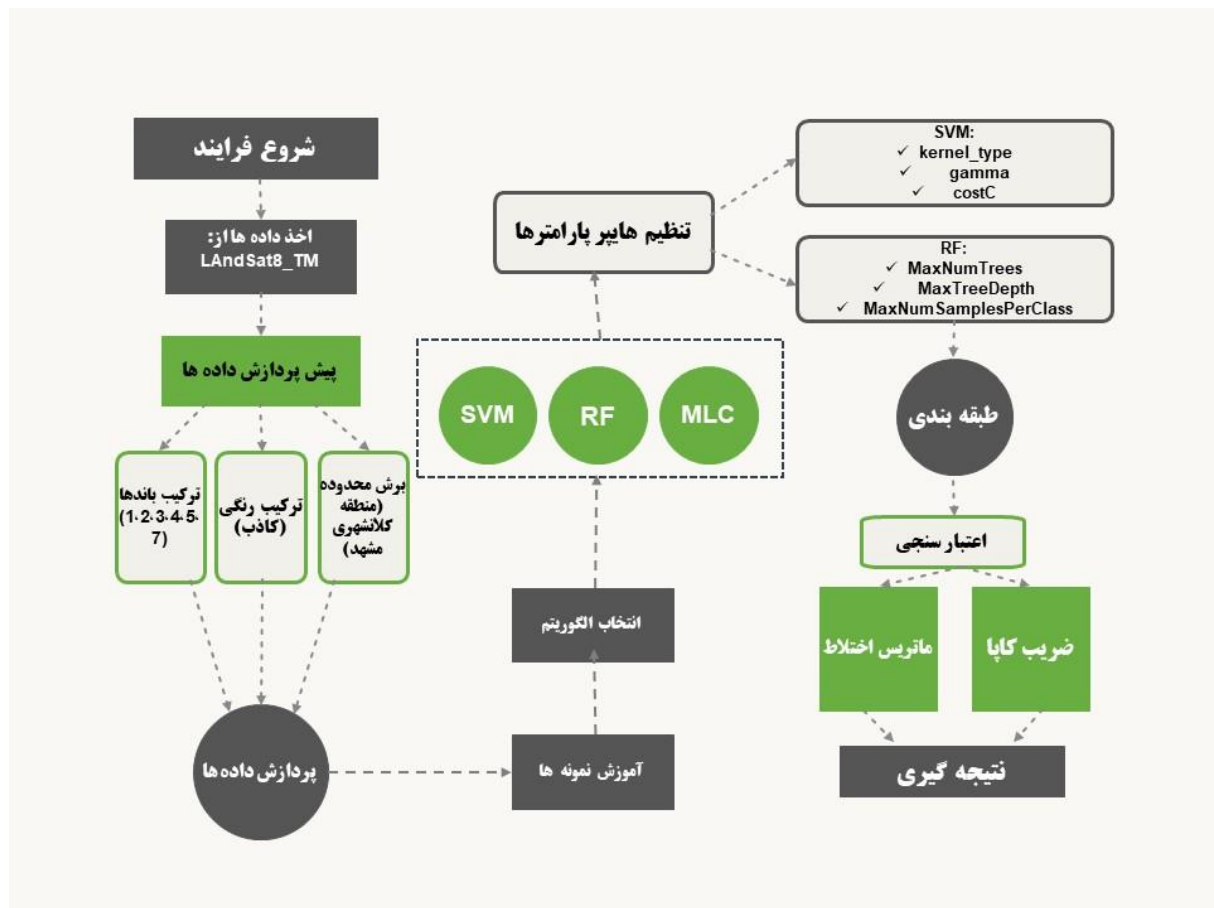
قاعدتاً طبقه بندى در محدوده اى ایستا و نسبتاً پایدار آسان تر از منظره هاى است که مانند منطقه کلان شهرى مشهد از وسعت، تنوع و پویایى این چنین بالایى برخوردار است، لذا عملکرد قابل قبول یک الگوریتم در طبقه بندى چنین منطقه اى، قابلیت استفاده و تعمیم آن را اثبات مى کند.



شکل ۱. محدوده مورد مطالعه

۲-۲- روش تحقیق

روش تحقیق حاضر از منظر هدف، کاربردى و از منظر ماهیت، توصیفى-تحلیلى است. گردآورى اطلاعات در این پژوهش به روش اسنادى-کتابخانه اى انجام شده است. شکل ۲ نمودار جریان فرایند تحقیق را نشان مى دهد.



شکل ۲. فرایند تحقیق

۲-۲-۱- جمع آوری داده ها

در این مطالعه تصویر سنجنده OLI (Operational Land Imager) در ماهواره لندست ۸ (تاریخ ۲۰۲۰/۰۷/۰۸) فاقد ابر و غبار از پایگاه داده USGS Earth Explorer تهیه شده است (earthexplorer.usgs.gov). این تصویر به طور خاص به دلیل عدم وجود ابر و غبار از پایگاه داده USGS Earth Explorer (earthexplorer.usgs.gov) تهیه گردیده است. تاریخ انتخاب شده مصادف است با اوایل تابستان که به منظور افزایش دقت در شناسایی تفاوت ها و تمایز بین مناطق با پوشش گیاهی و مناطق اطراف آن در نظر گرفته شده است. در این زمان، اراضی دارای پوشش گیاهی علاوه بر این که فاقد ابر هستند در متراکم ترین حالت خود نیز قرار دارند که به کمک این ویژگی، امکان تحلیل دقیق تری از کاربری های اراضی فراهم می شود و کیفیت تصویر نیز به خاطر خلوص آن بالا می رود.

تصاویر ماهواره لندست به دلیل دوره تاریخی طولانی تر نسبت به ماهواره سنتینل انتخاب شده است؛ یکی از اهداف کلی طبقه بندی کاربری اراضی، پایش تغییرات در طول زمان است و تحلیل روند تغییرات کاربری ها با تصاویر این ماهواره بهتر صورت می گیرد. به طور کلی، لندست ۸ معمولاً برای تحقیق و پروژه های نیازمند دقت بالاتر و داده های تاریخی مناسب تر است، در حالی که سنتینل به عنوان یک گزینه بهتر برای کاربردهای نیازمند اسکن های تکراری سریع تر و همچنین نظارت و مدیریت منابع طبیعی مناسب است. لذا در این پژوهش از تصویر ماهواره لندست به منظور ارزیابی دقت عملکرد الگوریتم های طبقه بندی استفاده شده است.

۲-۲-۲- پیش پردازش تصاویر ماهواره‌ای

در بخش پیش پردازش تصاویر، تصویر دریافت شده با ترکیب مجموعه باندهای ۱، ۲، ۳، ۴، ۵ و ۷ به دست آمد (باتوجه به این که باند ۶ حرارتی است، مورد استفاده قرار نگرفته است). و با توجه به این که تصویر لندست ۸ توسط سازمان زمین شناسی ایالات متحده آمریکا به صورت زمین مرجع شده ارائه شده است، نیازی به تصحیح هندسی ندارند. در ادامه، برای تشخیص عوارض در تصاویر از ترکیب رنگی کاذب برای نمایش تصاویر استفاده شده است. به منظور نمایش ترکیب رنگی کاذب مادون قرمز نزدیک^۱، باند ۴ (مادون قرمز نزدیک)، بصورت کانال قرمز، باند ۳ (قرمز)، بصورت کانال سبز و باند ۲ (سبز)، بصورت کانال آبی نمایش داده شده است. سپس، باتوجه به این که وسعت تصاویر ماهواره‌ای بیشتر از منطقه کلان‌شهری مشهد بود، منطقه مورد مطالعه برش زده شده و سپس مورد پردازش قرار گرفته است.

۲-۲-۳- پردازش تصاویر ماهواره‌ای

طبقه‌بندی تصاویر ماهواره‌ای شاید به عنوان اصلی ترین مرحله پردازش محسوب شود که از این طریق امکان تبدیل فضای تصویر (بازتابش‌های ثبت شده در باندهای مختلف) به فضای واقعی (نقشه‌های پوشش زمین و کاربری اراضی) ممکن می‌شود. به جداسازی مجموعه‌های طیفی مشابه و تقسیم‌بندی طیف‌هایی که دارای رفتار یکسانی هستند طبقه‌بندی اطلاعات ماهواره‌ای گفته می‌شود. در این فرایند یک برچسب خاص به هر یک از پیکسل‌های تصاویر داده می‌شود و پیکسل‌های مشابه در یک طبقه و با یک برچسب مشترک قرار می‌گیرند. به طور کلی، الگوریتم‌های طبقه‌بندی تصاویر ماهواره‌ای به چهار دسته ۱- نظارت شده^۲ ۲- نیمه نظارت شده^۳ ۳- بدون نظارت^۴ و ۴- تقوینی^۵ طبقه‌بندی می‌شوند (Glielmo et al., ۲۰۲۱; Dahlke et al., ۲۰۲۰; Azizol Ismail, ۲۰۲۰). یادگیری تحت نظارت، فرایند آموزش مدل‌های یادگیری ماشین با استفاده از داده‌های برچسب‌گذاری شده است تا بتوانند پیش‌بینی‌ها یا دسته‌بندی‌های دقیقی انجام دهند (Dahlke et al., ۲۰۲۰). الگوریتم‌های ماشین بردار پشتیبان و جنگل تصادفی که در این پژوهش استفاده شده‌اند در گروه یادگیری نظارت شده قرار دارند. ماشین بردار پشتیبان، با پیدا کردن یک ابر صفحه^۶ که کلاس‌ها را به بهترین شکل از هم جدا می‌کند، کار می‌کند (Talukdar et al., ۲۰۲۰). (Santarsiero et al., ۲۰۲۲; Chandra & Bedi, ۲۰۲۱). اما، الگوریتم جنگل تصادفی، روشی برای مدل‌سازی تصمیم‌گیری است که با تقسیم داده‌ها به زیرمجموعه‌ها بر اساس ویژگی‌ها، درختی از سؤالات باینری ایجاد می‌کند تا به پیش‌بینی یا دسته‌بندی برسد (Moradi & Kasani, ۲۰۲۴; Marudi et al., ۲۰۲۴; Naswin & Wibowo, ۲۰۲۳; Jiao et al., ۲۰۲۰). الگوریتم حداکثر احتمال به عنوان یک الگوریتم تحت نظارت، روشی است برای برآورد پارامترهای مدل آماری که با انتخاب مقادیر پارامترها، احتمال مشاهده داده‌های موجود را به حداکثر می‌رساند (Ganesh et al., ۲۰۲۳; Yimer et al., ۲۰۲۳). الگوریتم‌های یادگیری بدون نظارت، فرایند آموزش مدل‌های یادگیری ماشین با استفاده از داده‌های بدون برچسب است تا الگوها و ساختارهای پنهان در داده‌ها شناسایی شوند (Sakai et al., ۲۰۱۷; Dahlke et al., ۲۰۲۰; Gan et al., ۲۰۱۳). الگوریتم‌های ماشین بردار پشتیبان، جنگل تصادفی و بیشترین احتمال به دلایل زیر در این پژوهش انتخاب شده‌اند.

^۱ NIR

^۲ Supervised Learning

^۳ Semi-supervised Learning

^۴ Unsupervised Learning

^۵ Reinforcement Learning

^۶ hyperplane

-تنوع در رویکردها: این سه الگوریتم نمایانگر رویکردهای مختلف در طبقه‌بندی تصاویر ماهواره‌ای هستند. ماشین بردار پشتیبان به‌عنوان یک الگوریتم مبتنی بر حاشیه، در تلاش است تا بهترین مرز تفکیک‌کننده را بین کلاس‌ها پیدا کند. جنگل تصادفی به‌عنوان یک الگوریتم مبتنی بر مجموعه، از چندین درخت تصمیم برای بهبود دقت و کاهش احتمال اور فیتینگ استفاده می‌کند و الگوریتم بیشترین احتمال به‌عنوان یک روش آماری، بر اساس توزیع‌های احتمالی عمل می‌کند. این تنوع مقایسه عملکرد الگوریتم‌ها را در شرایط مختلف ممکن می‌سازد (Talukdar et al., ۲۰۲۰; Marudi et al., ۲۰۲۴; Ganesh et al., ۲۰۲۳).

-کارایی در مسائل مختلف: این الگوریتم‌ها به طور گسترده‌ای در مسائل مختلف طبقه‌بندی و پیش‌بینی مورد استفاده قرار گرفته‌اند و در بسیاری از مطالعات نشان داده‌اند که عملکرد خوبی دارند. انتخاب این الگوریتم‌ها دستیابی به نتایج قابل مقایسه و معتبرتری را فراهم می‌آورد.

-سازگاری با داده‌های پیچیده: جنگل تصادفی و ماشین بردار پشتیبان به‌خوبی می‌توانند با داده‌های پیچیده و غیرخطی کار کنند. این ویژگی‌ها به‌ویژه در زمینه‌های جغرافیایی و محیطی که داده‌ها ممکن است دارای الگوهای پیچیده باشند، اهمیت دارد.

-قابلیت تفسیر: الگوریتم‌های مورد مطالعه، قابلیت تفسیر بالایی دارند و می‌تواند به راحتی به کاربران کمک کند تا بفهمند که چگونه تصمیمات گرفته می‌شوند. این ویژگی می‌تواند در تحلیل نتایج و ارائه آن‌ها به ذی‌نفعان مفید باشد.

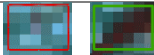
-تجربه‌های پیشین: این الگوریتم‌ها در بسیاری از پژوهش‌های پیشین به‌عنوان الگوریتم‌های مؤثر در زمینه‌های مشابه شناخته شده‌اند. این امر می‌تواند به اعتبار و قابلیت اعتماد نتایج کمک نماید.

در هر الگوریتم طبقه‌بندی، معایب خاصی نیز وجود دارد، بنابراین انتخاب الگوریتم‌ها باید با دقت صورت گیرد و در این فرایند، ویژگی‌های تحقیق و محدوده مورد مطالعه باید مد نظر قرار گیرد. الگوریتم ماشین بردار پشتیبان ممکن است در مواجهه با داده‌های بزرگ و با ابعاد بالا با مشکلاتی در زمان آموزش و پیش‌بینی روبه‌رو شود و پیچیدگی محاسباتی بالایی را به همراه داشته باشد (Talukdar et al., ۲۰۲۰; Santarsiero et al., ۲۰۲۲; Chandra & Bedi, ۲۰۲۱). همچنین، جنگل تصادفی ممکن است در کار با داده‌های نامتعادل کارایی مطلوبی نداشته باشد و نیاز به مصرف بالای حافظه داشته باشد (Piryonesi, ۲۰۱۹; Louppe, ۲۰۱۴). با این حال، با توجه به محدودیت‌های وسعت محدوده مورد مطالعه در این پژوهش، معایب مذکور در اینجا بروز نخواهند کرد. علاوه بر این، الگوریتم بیشترین احتمال فرض‌های ساده‌انگاری زیادی در مورد استقلال ویژگی‌ها ایجاد می‌کند که ممکن است در شرایط واقعی نادرست باشد، هرچند که این ویژگی معمولاً در داده‌های جغرافیایی برآورد نمی‌شود (Ganesh et al., ۲۰۲۳; Yimer et al., ۲۰۲۲). به همین دلیل، الگوریتم‌های انتخاب شده گزینه‌های مناسبی برای طبقه‌بندی منطقه کلان‌شهری مشهد به شمار می‌آیند.

هریک از الگوریتم‌های طبقه‌بندی دارای معایبی نیز هستند، لذا انتخاب الگوریتم‌ها می‌بایست با احتیاط صورت گیرد و در این فرایند، ویژگی‌های تحقیق و محدوده مورد مطالعه مد نظر قرار گیرد. الگوریتم ماشین بردار پشتیبان ممکن است با داده‌های بزرگ و با ابعاد بالا به راحتی دچار مشکلات زمان آموزش و پیش‌بینی شوند و دارای پیچیدگی محاسباتی شوند (Talukdar et al., ۲۰۲۰; Santarsiero et al., ۲۰۲۲; Chandra & Bedi, ۲۰۲۱). جنگل تصادفی نیز می‌تواند با داده‌های نامتعادل به درستی عمل نکند و همچنین نیاز به مصرف حافظه بالایی دارد (Piryonesi, ۲۰۱۹; Louppe, ۲۰۱۴) اما با توجه به این که وسعت محدوده مورد مطالعه در این پژوهش زیاد نیست، لذا معایب مذکور در این پژوهش بروز پیدا نمی‌کنند. همچنین، الگوریتم بیشترین احتمال فرض‌های سادگی زیادی را در مورد استقلال ویژگی‌ها ایجاد می‌کند که در موارد واقعی ممکن است نادرست باشد، اما این ویژگی معمولاً در داده‌های جغرافیایی برآورد نمی‌شود (Ganesh et al., ۲۰۲۳; Yimer et al., ۲۰۲۲). لذا، الگوریتم‌های منتخب گزینه‌های مناسبی برای طبقه‌بندی منطقه کلان‌شهری مشهد به شمار می‌آیند.

مرحله بعدی، تهیه نمونه‌های تعلیمی از واقعیت زمینی است. این نمونه‌ها به دو گروه تقسیم می‌شوند: ۱- گروه آموزش که تعداد ۱۴۵۰ نمونه برای کلیه طبقات انتخاب شده است (۰.۷۵٪) و ۲- گروه آزمایش که حدود تعداد ۵۰۰ نقطه به عنوان نمونه آموزشی، وارد ماشین گشته برچسب زده شده و آموزش می‌بینند (۰.۲۵٪). باتوجه به دقت تصاویر ماهواره‌های لندست ۸ (۳۰ متر) که از سال ۲۰۲۰ تهیه شده است و ویژگی‌های منطقه کلان‌شهری مشهد، نمونه‌های منتخب در ۵ طبقه ۱- مناطق ساخته شده^۱ ۲- اراضی بایر^۲ ۳- مناطق کوهستانی^۳ ۴- فضاهای سبز^۴ و ۵- پهنه‌های آبی^۵ دسته‌بندی شده‌اند.

جدول ۱. کلاس‌های طبقه‌بندی

نام طبقه	توضیح	مثالی از نمونه انتخابی طبقه
مناطق ساخته شده	مناطق مسکونی، تجاری، صنعتی، شبکه راه‌های و هر نوع تأسیسات و تجهیزات شهری	
ارضای بایر	ارضای با پوشش خاک و ماسه	
مناطق کوهستانی	مناطق کوهستانی از جنس سنگ و صخره	
فضاهای سبز	پارک‌ها، درختان، علفزارها، مناطق جنگلی، اراضی کشاورزی با انواع محصولات	
پهنه‌های آبی	سد، باتلاق، رودخانه، کانال، استخرهای روباز	

۲-۲-۴- تنظیم هایپرپارامترها

در این بخش از پژوهش به تنظیم هایپرپارامترهای مدل پرداخته شده است. هایپرپارامترها مقادیر قابل تنظیمی هستند که قبل از آموزش مدل‌های یادگیری ماشین تعیین می‌شوند و بر عملکرد و دقت مدل تأثیر می‌گذارند. برخلاف پارامترهای داخلی مدل که در خلال فرآیند آموزش به دست می‌آیند، هایپرپارامترها باید به صورت دستی یا با استفاده از تکنیک‌های بهینه‌سازی تنظیم شوند. تنظیم مناسب هایپرپارامترها می‌تواند به بهبود قابلیت تعمیم و کارایی مدل کمک کند (Alonso-Sarría et al, ۲۰۲۴). در این پژوهش پس از انتخاب نمونه‌ها، هایپرپارامترها با تکنیک بهینه‌سازی تنظیم شده است؛ این بدین معنا است که نرم‌افزار با توجه به تعداد نمونه‌ها و پیچیدگی آن‌ها، مقادیر هایپرپارامترها را به صورت خودکار تنظیم می‌نماید. این مقادیر می‌تواند پس از صحت‌سنجی مدل مورد ارزیابی قرار گیرد و به صورت دستی تغییر کند.

در الگوریتم ماشین بردار پشتیبان، از تابع کرنل گوسی^۶ به همراه هایپرپارامترهای مقدار گاما ۳۲ و C برابر با ۸۱۹۲ استفاده شده است. مقدار بزرگ گاما (۳۲) به این معناست که تابع کرنل به دقت بیشتری به داده‌های آموزشی نزدیک خواهد شد و قادر است به پیچیدگی‌های موجود در توزیع داده‌ها پاسخ دهد؛ این امر برای طبقه‌بندی دقیق کاربری اراضی که ممکن است الگوهای بسیار متفاوتی داشته باشند، ضروری است. از سوی دیگر، مقدار C برابر با ۸۱۹۲ نشان‌دهنده تمایل به حداقل نگه‌داشتن خطا در داده‌های آموزشی بوده و نشان‌دهنده یک مرز تفکیکی قوی و دقیق است. این امر به موازات قابلیت مدل، در تفکیک دقیق

^۱ Built up areas

^۲ Barren lands

^۳ Rocks

^۴ Green spaces

^۵ Water bodies

^۶ RBF

کلاس‌ها کمک می‌کند؛ به طور کلی، این ترکیب از هایپرپارامترها در راستای به حداکثر رساندن کارایی و دقت طبقه‌بندی مدل در یک منطقه پیچیده، مانند مشهد، تنظیم شده‌اند.

در الگوریتم جنگل تصادفی، حداکثر تعداد درختان^۱ برابر با ۵۰ و حداکثر عمق آن‌ها^۲ برابر با ۳۰ تنظیم شده است. حداکثر تعداد نمونه‌ها برای هر کلاس نیز برابر با ۱۰۰۰ در نظر گرفته شده است. این مقادیر از هایپرپارامترها به دلیل ایجاد توازن مناسب میان دقت و کارایی مدل در تحلیل داده‌های جغرافیایی انجام شده است. تعداد ۵۰ درخت، تنوع کافی را ارائه می‌دهد تا بتواند ویژگی‌های مختلف داده را به خوبی شناسایی کند و از خطر بیش‌برازش جلوگیری کند. عمق حداکثر ۳۰ برای درختان، توانایی مدل را در یادگیری نازک‌کاری‌های پیچیده‌تری که در داده‌های جغرافیایی ممکن است وجود داشته باشد، افزایش می‌دهد. همچنین، محدود کردن حداکثر تعداد نمونه‌ها به ۱۰۰۰ در هر کلاس، به بهینه‌سازی اندازه هر زیرمجموعه آموزشی کمک می‌کند و از نامتعادل شدن داده‌ها جلوگیری می‌کند. این ترکیب از هایپرپارامترها نشان‌دهنده رویکردی دقیق در طراحی مدل برای شناسایی الگوهای جغرافیایی است.

پس از انتخاب نمونه، برچسب‌زدن آن‌ها در قالب طبقات تعریف شده و تنظیم هایپرپارامترها، مرحله اصلی مدل، یعنی پیش‌بینی سایر پیکسل‌های برچسب زده نشده در قالب طبقات تعریف شده در الگوریتم‌های ماشین بردار پشتیبان، جنگل تصادفی و بیشترین احتمال آغاز می‌شود. در این مرحله پس از وارد کردن تصویر ماهواره‌ای موردنظر و نمونه‌های برداری نقطه‌ای به ماشین، تصویر طبقه‌بندی شده کاربری اراضی به عنوان خروجی مدل در نظر گرفته می‌شود.

۲-۲-۵- ارزیابی صحت طبقه‌بندی مدل

در راستای ارزیابی دقت مدل، می‌بایست عملکرد مدل ایجاد شده در برابر مجموعه داده‌هایی که قبلاً وارد مدل نگشته است را بررسی نمود. در واقع، یک ماشین، الگوها و خصوصیات مختلفی را از داده‌هایی که به آن آموزش داده شده است یاد می‌گیرد و خود را برای تصمیم‌گیری‌هایی در زمینه‌های مختلف مانند شناسایی، طبقه‌بندی یا پیش‌بینی داده‌های جدید آموزش می‌دهد. برای بررسی دقیق این که ماشین چگونه قادر به اتخاذ این تصمیمات است، پیش‌بینی‌ها را بر روی داده‌های آموزش داده شده، آزمایش می‌کنند. برای این کار ابتدا بر روی داده‌های آموزش داده شده کار می‌کنیم و پس از آموزش مدل به اندازه کافی، از آن برای آزمایش بر اساس داده‌ها استفاده می‌کنیم.

برای ارزیابی صحت طبقه‌بندی دو روش وجود دارد:

۱- برآورد ماتریس اختلاط^۳ (خطا)

۲- محاسبه ضریب توافق کاپا^۴

۱- برآورد ماتریس اختلاط

این روش یکی از متداول‌ترین روش‌های ارزیابی صحت طبقه‌بندی است که رابطه بین داده‌های مرجع (حقایق زمینی) و نتایج طبقه‌بندی خودکار را مقایسه می‌نماید. در این مرحله نقاط آزمایشی طبقه‌بندی شده را با تصویر ماهواره‌ای و تصاویر گوگل ارث مقایسه نموده و طبقه مدل‌سازی شده را با کاربری واقعی مقایسه نموده و سپس صحت یا عدم صحت طبقه نقطه آزمایشی را

^۱ MaxNumTrees

^۲ MaxTreeDepth

^۳ Confusion Matrix

^۴ Cohen's kappa coefficient

مشخص می‌نماییم. این ماتریس با یک ساختار جدولی نشان می‌دهد که چه تعداد از رده‌ها به اشتباه به جای رده دیگری رده‌بندی شده‌اند. در یادگیری بدون نظارت به آن ماتریس تطابق^۱ هم گفته می‌شود.

در هر طبقه‌بندی برای هر نمونه داده، یکی از چهار حالتی که در ادامه بیان شده، ممکن است اتفاق بیفتد

نمونه عضو دسته مثبت باشد و عضو همین کلاس تشخیص داده شود (مثبت صحیح)^۲

نمونه عضو کلاس مثبت باشد و عضو کلاس منفی تشخیص داده شود (منفی کاذب)^۳

نمونه عضو کلاس منفی باشد و عضو همین کلاس تشخیص داده شود (منفی صحیح)^۴

نمونه عضو کلاس منفی باشد و عضو کلاس مثبت تشخیص داده شود (مثبت کاذب)^۵

پس از اجرای الگوریتم دسته‌بندی، باتوجه به توضیحات و تعاریف ذکر شده، می‌توان عملکرد یک طبقه‌بند را به کمک جدول ۲ بررسی کرد (Arias-Duart et al., ۲۰۲۳).

جدول ۲. حالات مختلف سلول‌های طبقه‌بندی شده

		برچسب پیش‌بینی شده	
		مثبت	منفی
برچسب شناخته شده	مثبت	TP	FN
	منفی	FP	TN

این جدول را اصطلاحاً ماتریس درهم‌ریختگی می‌گویند. جدول یا ماتریس درهم‌ریختگی، نتایج حاصل از طبقه‌بندی را بر اساس اطلاعات واقعی موجود، نمایش می‌دهد. حال بر اساس این مقادیر می‌توان معیارهای مختلف ارزیابی و اندازه‌گیری دقت را تعریف کرد. معیار صحت^۶، متداول‌ترین، اساسی‌ترین و ساده‌ترین معیار اندازه‌گیری کیفیت یک دسته‌بند است و عبارت است از میزان تشخیص صحیح دسته‌بند در مجموع دودسته. این معیار در واقع نشانگر میزان الگوهای است که درست تشخیص داده شده‌اند و بر اساس ماتریس ارائه شده در بالا، به شکل زیر فرموله و تعریف می‌شود (Yinet al., ۲۰۱۹).

$$\text{Accuracy} = \frac{TP+TN}{TP+FN+FP+TN}$$

رابطه ۱. معیار صحت

۲- ضریب User's accuracy

این ضریب موارد مثبت کاذب^۷ را نشان می‌دهد، جایی که پیکسل‌ها به اشتباه به عنوان یک کلاس شناخته شده طبقه‌بندی می‌شوند درحالی‌که باید به عنوان چیز دیگری طبقه‌بندی می‌شدند. این ضریب به عنوان خطاهای کمیسیون یا خطای نوع اول نیز گفته می‌شود. داده‌های محاسبه این میزان خطا از ردیف‌های ماتریس اختلاط خوانده می‌شود.

۳- ضریب Producer's accuracy

^۱ Matching Matrix

^۲ True Positive

^۳ False Negative

^۴ True Negative

^۵ False Positive

^۶ Accuracy

^۷ False Positive (FP)

این ضریب موارد منفی کاذب^۱ است که در آن پیکسل‌های یک کلاس شناخته شده به عنوان چیزی غیر از آن کلاس طبقه‌بندی می‌شوند. این ضریب به عنوان خطاهای حذف یا خطای نوع دوم نیز شناخته می‌شود. داده‌های محاسبه این میزان خطا در ستون‌های جدول خوانده می‌شود.

۴- ضریب توافق کاپا^۲

ضریب کاپا تکنیک چندمتغیره گسسته‌ای است که اگر یک ماتریس خطا تفاوت معناداری با ماتریس دیگر داشته باشد، در ارزیابی دقت برای تصمیم‌گیری آماری مورد استفاده قرار می‌گیرد. این ضریب عددی بین ۱- تا ۱+ است که هر چه به ۱+ نزدیک‌تر باشد، بیانگر وجود توافق بیشتر بین مقیاس‌ها یا ارزیاب‌ها است و هر چه به ۱- نزدیک‌تر باشد، نشان‌دهنده وجود توافق کمتر بین آن‌ها است. از طرفی اگر ضریب توافق کاپا صفر شود، نشان‌دهنده عدم توافق کامل است (Nti et al., ۲۰۲۳). جدول شماره ۳ نشان‌دهنده دستورالعمل‌هایی است که توسط لندیس و کوچ، فلیس و آلمن برای ارزیابی سطح توافق ارائه شده است (Landis & Koch, ۱۹۹۷; Fleiss et al., ۱۹۷۱; Altman, ۱۹۹۵).

جدول ۳. آستانه‌های قدرت ضریب توافق کاپا

مقیاس معیار آلمن		مقیاس معیار فلیس		مقیاس معیار لندیس و کوچ	
ناسازگار	<۰.۲	ناسازگار	<۰.۲	ضعیف	<۰.۴
ضعیف	۰.۲۱-۰.۴۰	ضعیف	۰.۲۱-۰.۴۰	متوسط تا خوب	۰.۴۰-۰.۷۵
نسبتاً خوب	۰.۴۱-۰.۶۰	نسبتاً خوب	۰.۴۱-۰.۶۰	عالی	بیشتر از ۰.۷۵
خوب	۰.۶۱-۰.۸۰	محکم	۰.۶۱-۰.۸۰		
بسیار خوب	۰.۸۱-۱.۰۰	تقریباً کامل	۰.۸۱-۱.۰۰		

منبع: (Landis & Koch, ۱۹۹۷; Fleiss et al., ۱۹۷۱; Altman, ۱۹۹۵)

$$\text{kappa} = \text{Pi} = (\text{PA}_0 - \text{PA}_E) / (1 - \text{PA}_E)$$

رابطه ۲. ضریب کاپا

مقدار PA_0 نمایانگر میزان توافق دو ارزیاب است

مقدار PA_E نیز نمایانگر میزان توافق مورد انتظار است.

۴- یافته‌های تحقیق

۴-۱- طبقه بندی کاربری اراضی

تحلیل فضایی کاربری اراضی در منطقه کلان‌شهری مشهد نشان می‌دهد که پوشش اراضی بایر و پهنه‌های آبی تقریباً به یک اندازه در همه طبقه‌بندی‌کننده‌ها طبقه‌بندی می‌شوند، به گونه‌ای که در هر سه الگوریتم اراضی بایر با پوشش بیش از ۴۰ درصد از محدوده، کاربری غالب در هر سه الگوی طبقه بندی در این منطقه به شمار می‌رود و کاربری پهنه‌های آبی نیز که در هر سه الگوریتم کم‌تر از یک درصد از محدوده را تشکیل می‌دهد، کم‌ترین کاربری موجود در این ناحیه به شمار می‌آید. جدول ۴ درصد سهم هر کلاس را با توجه به کل پوشش زمین در منطقه مورد مطالعه برای هر طبقه‌بندی نشان می‌دهد. با توجه به جدول ۴ و

^۱ False Negative (FN)

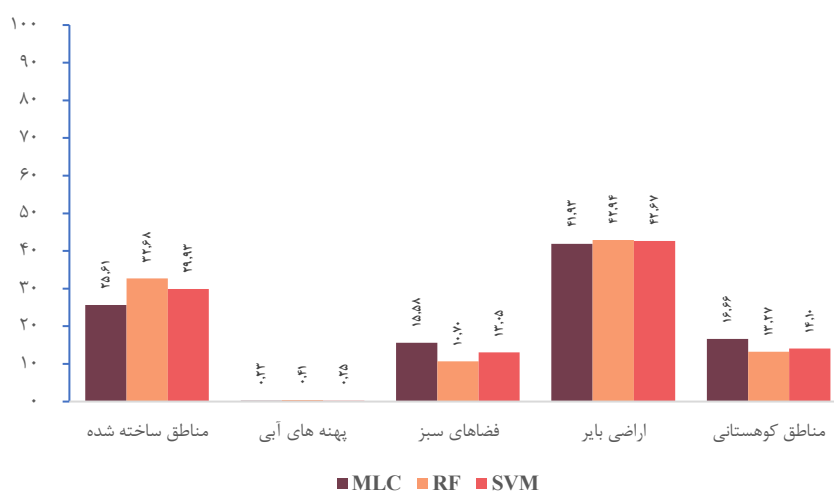
^۲ Kappa

شکل ۳. مساحت مناطق ساخته شده در الگوریتم‌های مورد بررسی تفاوت زیادی را بروز داده‌اند به گونه‌ای که این طبقه در الگوریتم جنگل تصادفی بیش‌ترین میزان (۳۲/۶۸٪) و در الگوریتم بیش‌ترین احتمال کمترین میزان (۲۵/۶۱٪) را نشان می‌دهد. فضاهای سبز نیز در سه الگوریتم تفاوت زیادی را نشان می‌دهند. اما بر خلاف طبقه مناطق ساخته شده، این طبقه در الگوریتم بیش‌ترین احتمال، بیش‌ترین میزان (۱۵/۵۸٪) و در الگوریتم جنگل تصادفی کمترین میزان (۱۰/۷۰٪) را دارا می‌باشد.

جدول ۴. سهم طبقات مختلف کاربری اراضی در منطقه کلان‌شهری مشهد در سال ۲۰۲۰

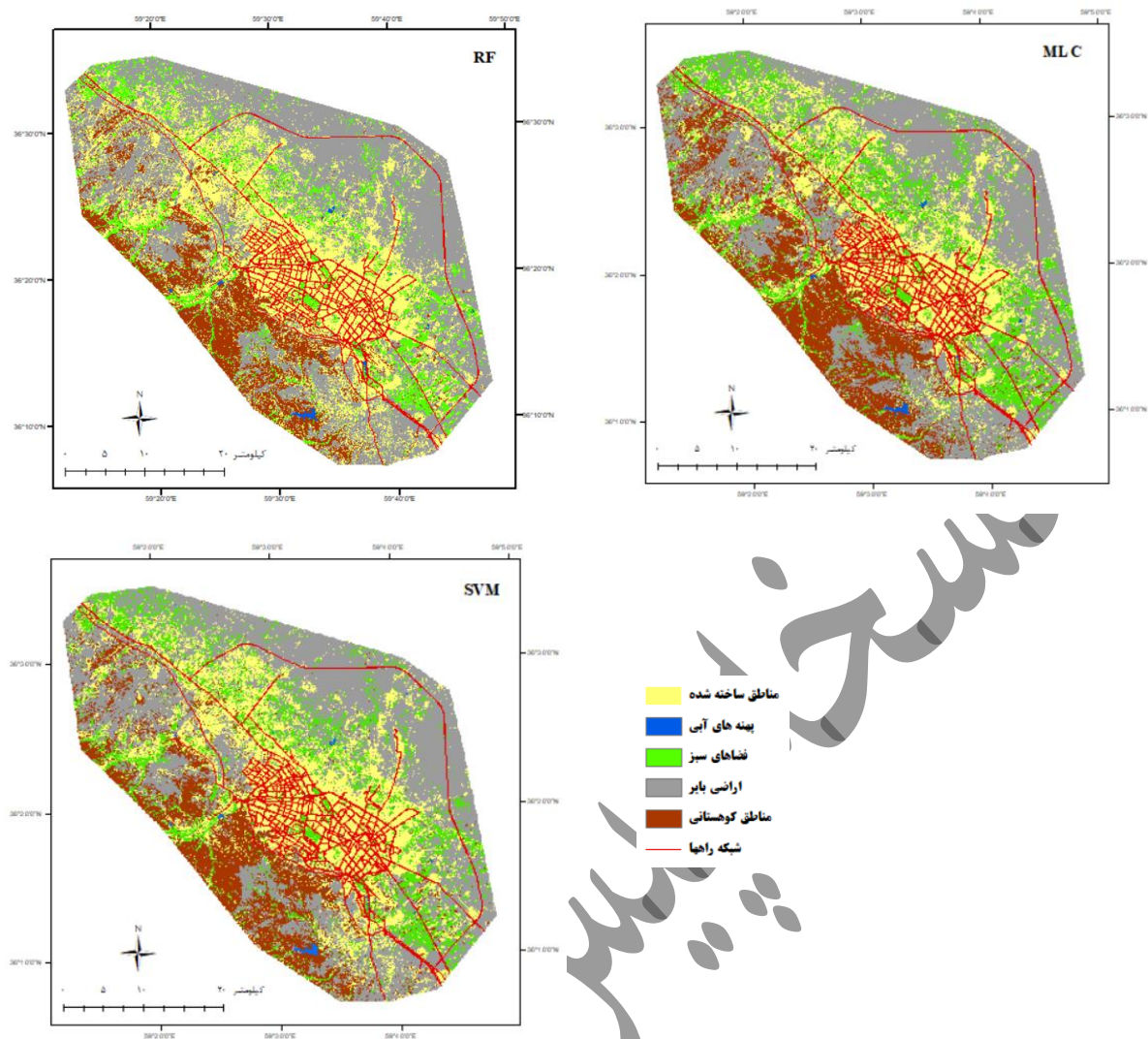
طبقات کاربری	مناطق ساخته شده		پهنه‌های آبی		فضاهای سبز		اراضی بایر		مناطق کوهستانی		مساحت کل (هکتار)
	هکتار	درصد	هکتار	درصد	هکتار	درصد	هکتار	درصد	هکتار	درصد	
بیشترین احتمال (MLC)	۴۶۴۵۶/۸۵	۲۵/۶۱	۴۱۷/۵۱	۰/۲۳	۲۸۲۶۵/۱۲	۱۵/۵۸	۷۶۰۵۶/۷۷	۴۱/۹۳	۳۰۲۱۴/۱۴	۱۶/۶۶	۱۸۱۴۱۰/۴۰
جنگل تصادفی (RF)	۵۹۲۸۰/۹۳	۳۲/۶۸	۷۳۹/۷۸	۰/۴۱	۱۹۴۱۴/۴۶	۱۰/۷۰	۷۷۸۹۸/۲۱	۴۲/۹۴	۲۴۰۷۶/۶۳	۱۳/۲۷	
ماشین بردار پشتیبان (SVM)	۵۴۳۰۰/۰۰	۲۹/۹۳	۴۵۱/۰۶	۰/۵	۲۳۶۶۸/۳۹	۱۳/۰۵	۷۷۴۰۷/۸۷	۴۲/۶۷	۲۵۵۸۳/۴۲	۱۴/۱۰	
انحراف معیار (STD)	-	۳/۵۶	-	۰/۰۹	-	۲/۴۴	-	۰/۵۲	-	۱/۷۶	
میانگین (AVE)	-	۲۹/۴۱	-	۰/۳۰	-	۱۳/۱۱	-	۴۲/۵۱	-	۱۴/۶۸	
ضریب تغییرات (CV)(٪)	-	۱۲/۱۱	-	۳۳/۰۴	-	۱۸/۶۱	-	۱/۲۳	-	۱۲/۰۱	

منبع: یافته‌های تحقیق



شکل ۳. سهم طبقات مختلف کاربری اراضی در منطقه کلان‌شهری مشهد در سال ۲۰۲۰

نقشه‌های کاربری اراضی منطقه کلان‌شهری مشهد از طریق جنگل تصادفی، ماشین بردار پشتیبان و بیش‌ترین احتمال در شکل ۴ آمده است. مقایسه توزیع فضایی کاربری‌های طبقه‌بندی شده نشان می‌دهد که طبقه مناطق ساخته شده در الگوریتم جنگل تصادفی، بیش از سایر الگوریتم‌ها در میان مناطق کوهستانی و فضاهای سبز پیش رفته است.



شکل ۴. طبقه بندی کاربری اراضی منطقه کلان شهری مشهد با الگوریتم‌های جنگل تصادفی (RF)، ماشین بردار پشتیبان (SVM) و بیشترین شباهت (MLC)

انحراف معیار (SD) و ضریب تغییرات (CV) درصد سهم مساحت در یک کلاس LULC توسط الگوریتم‌های مختلف نیز در جدول ۴ نشان داده شده است. ضرایب محاسبه شده به وضوح نشان می‌دهند که اراضی بایر بادقت بیشتری طبقه‌بندی می‌شوند، زیرا همه طبقه‌بندی‌کننده‌ها دارای مناطق کاملاً یکنواخت با ضریب تغییرات بسیار کم هستند. در مقابل، پهنه‌های آبی و فضاهای سبز به خوبی طبقه‌بندی نشده‌اند، زیرا همه الگوریتم‌ها سهم وسعت طبقات مذکور را با تفاوت‌های قابل توجهی در نظر گرفته‌اند.

۴-۲-۴ اعتبارسنجی طبقه بندی‌ها

دقت کلی هر یک از الگوریتم‌ها با استفاده از ماتریس اختلاط و ضریب کاپا (K) در جدول ۵ تا ۷ نشان داده است. طبقه‌بندی‌کننده SVM به عنوان مدل بسیار دقیق LULC با ضریب کاپا ۰.۹۳ در بین همه طبقه‌بندی‌کننده‌ها شناسایی شده است. لذا این الگوریتم بین نقشه طبقه‌بندی‌شده LULC و واقعیت زمینی تطابق بسیار خوبی را نشان می‌دهند. ضریب کاپا در مدل RF برابر است با ۰.۸۸ و در مدل MLC برابر است با ۰.۸. بنابراین همه مدل‌ها را می‌توان مفید دانست، اما الگوریتم SVM را می‌توان به عنوان بهترین طبقه بندی‌کننده مناسب LULC توصیه کرد.

جدول ۵. ماتریس اختلاط کاربری اراضی منطقه کلان‌شهری مشهد در الگوریتم جنگل تصادفی (RF)

کلاس کاربری	مناطق ساخته شده	پهنه‌های آبی	فضاهای سبز	اراضی بایر	مناطق کوهستانی	مجموع	U_Accuracy	Kappa
مناطق ساخته شده	۱۲۵	۰	۵	۸	۲۶	۱۶۴	۰.۷۶۲۱۹۵۱۲۲	۰
پهنه‌های آبی	۱	۷	۲	۰	۰	۱۰	۰.۷	۰
فضاهای سبز	۰	۰	۵۴	۰	۰	۵۴	۱	۰
اراضی بایر	۰	۰	۰	۲۱۴	۰	۲۱۴	۱	۰
مناطق کوهستانی	۰	۰	۰	۰	۶۶	۶۶	۱	۰
مجموع	۱۲۶	۷	۶۱	۲۲۲	۹۲	۵۰۸	۰	۰
P_Accuracy	۰.۹۹۲۰۶۳۵	۱	۰.۸۸۵۲۴۵۹	۰.۹۶۳۹۶ ۴	۰.۷۱۷۳۹۱۳	۰	۰.۹۱۷۳۲۲۸۳۵	۰
Kappa	۰	۰	۰	۰	۰	۰	۰	۰.۸۸۱۷۶۶

منبع: یافته‌های تحقیق

بررسی ضریب U_Accuracy در الگوریتم جنگل تصادفی نشان می‌دهد که برخی سلول‌ها به اشتباه به عنوان مناطق ساخته شده برداشت شده‌اند. یعنی در طبقه‌بندی این سلول‌ها خطای نوع اول رخ داده و سلول‌هایی که در طبقه مناطق ساخته شده قرار نمی‌گیرند به اشتباه در این طبقه قرار گرفتند (میزان صحت = 0.762195122). خطای نوع اول برای طبقه پهنه‌های آبی نیز مشاهده می‌شود. یعنی سلول‌هایی که آب نیستند اما به اشتباه آب تشخیص داده شده‌اند (میزان صحت = 0.7). از طرفی بررسی ضریب P_Accuracy نشان می‌دهد که خطای نوع دوم در طبقات مناطق کوهستانی (میزان صحت = 0.7173913) و فضاهای سبز (میزان صحت = 0.8852459) بیش از سایر طبقات است به گونه‌ای که سلول‌های این دو طبقه به درستی تشخیص داده نشده و به عنوان کاربری دیگری طبقه بندی شده است.

در مجموع کاربری بایر با کم‌ترین خطای نوع اول و دوم نسبت به سایر کاربری‌ها شناسایی شده است (میزان صحت به ترتیب برابر است با ۱ و 0.963964). اما به طور کلی با بررسی خطاهای نوع اول و دوم می‌توان گفت که درصد صحت این خطاها از بازه قابل قبول تجاوز نمی‌نماید.

جدول ۶. ماتریس اختلاط کاربری اراضی منطقه کلان‌شهری مشهد در الگوریتم ماشین بردار پشتیبان (SVM)

کلاس کاربری	مناطق ساخته شده	پهنه‌های آبی	فضاهای سبز	اراضی بایر	مناطق کوهستانی	مجموع	U_Accuracy	Kappa
مناطق ساخته شده	۱۴۱	۰	۲	۶	۰	۱۴۹	۰.۹۴۶۳۰۸۷	۰
پهنه‌های آبی	۱	۹	۰	۰	۰	۱۰	۰.۹	۰
فضاهای سبز	۰	۰	۶۶	۰	۰	۶۶	۱	۰
اراضی بایر	۲	۰	۰	۲۱۱	۰	۲۱۳	۰.۹۹۰۶۱۰۳	۰
مناطق کوهستانی	۴	۰	۰	۷	۶۰	۷۱	۰.۸۴۵۰۷۰۴	۰
مجموع	۱۴۸	۹	۶۸	۲۲۴	۶۰	۵۰۹	۰	۰
P_Accuracy	۰.۹۵۲۷۰۲۷	۱	۰.۹۷۰۵۸۸۲	۰.۹۴۱۹۶۴ ۳	۱	۰	۰.۹۵۶۷۷۸	۰
Kappa	۰	۰	۰	۰	۰	۰	۰	۰.۹۳۷۹۵

منبع: یافته‌های تحقیق

بررسی ضریب $U_Accuracy$ در الگوریتم ماشین بردار پشتیبان نشان می‌دهد که تعداد سلول‌هایی که به اشتباه به عنوان مناطق کوهستانی برداشت شده‌اند بیش از سایر کاربری‌ها است، یعنی در طبقه‌بندی این سلول‌ها خطای نوع اول رخ داده و سلول‌هایی که در طبقه مناطق کوهستانی قرار نمی‌گیرند به اشتباه در این طبقه قرار گرفتند (میزان صحت = $0/8450704$). از طرفی بررسی ضریب $P_Accuracy$ نشان می‌دهد که خطای نوع دوم در طبقات مختلف به مقدار بسیار ناچیزی دیده می‌شود به گونه‌ای که حداقل این ضریب در اراضی بایر دیده می‌شود که برابر است با $0/9419643$. در مجموع خطای نوع اول و دوم در این الگوریتم بسیار ناچیز است.

جدول ۷. ماتریس اختلاط کاربری اراضی منطقه کلان‌شهری مشهد در الگوریتم بیشترین احتمال (MLC)

کلاس کاربری	مناطق ساخته شده	پهنه‌های آبی	فضاهای سبز	اراضی بایر	مناطق کوهستانی	مجموع	$U_Accuracy$	Kappa
مناطق ساخته شده	۱۱۱	۰	۷	۸	۲	۱۲۸	$0/8671875$	۰
پهنه‌های آبی	۱	۹	۰	۰	۰	۱۰	$0/9$	۰
فضاهای سبز	۰	۰	۷۹	۰	۰	۷۹	1	۰
اراضی بایر	۲۷	۰	۴	۱۷۸	۰	۲۰۹	$0/8516746$	۰
مناطق کوهستانی	۷	۰	۰	۱۴	۶۲	۸۳	$0/746988$	۰
مجموع	۱۴۶	۹	۹۰	۲۰۰	۶۴	۵۰۹	۰	۰
$P_Accuracy$	$0/760274$	۱	$0/8777778$	$0/89$	$0/96875$	۰	$0/8624754$	۰
Kappa	۰	۰	۰	۰	۰	۰	۰	$0/8085242$

منبع: یافته‌های تحقیق

بررسی ضریب $U_Accuracy$ در الگوریتم بیشترین احتمال نشان می‌دهد که برخی سلول‌ها به اشتباه به عنوان مناطق کوهستانی شناسایی شده‌اند. یعنی در طبقه‌بندی این سلول‌ها خطای نوع اول رخ داده و سلول‌هایی که در طبقه مناطق کوهستانی قرار نمی‌گیرند به اشتباه در این طبقه قرار گرفتند (میزان صحت = $0/746988$). خطای نوع اول برای طبقه مناطق ساخته شده (میزان صحت = $0/8671875$) و اراضی بایر (میزان صحت = $0/8516746$) نیز مشاهده می‌شود. از طرفی بررسی ضریب $P_Accuracy$ نشان می‌دهد که خطای نوع دوم به ترتیب در طبقات مناطق ساخته شده (میزان صحت = $0/760274$) و فضاهای سبز (میزان صحت = $0/8777778$) بیش از سایر طبقات است به گونه‌ای که سلول‌های این طبقه به درستی تشخیص داده نشده و به عنوان کاربری دیگری طبقه‌بندی شده است. با این وجود، با بررسی خطاهای نوع اول و دوم می‌توان گفت که درصد صحت این خطاها از بازه قابل قبول تجاوز نمی‌نماید.

۵- بحث

پژوهش‌ها و مطالعات مختلف نشان داده است که الگوریتم‌های مختلف در شرایط متفاوت دقت متفاوتی دارند. فرانسیسکو آلونسو ساریا (۲۰۲۴)، به این نتیجه رسیده است که جنگل تصادفی (RF) و ماشین بردار پشتیبان (SVM) تا حدودی عملکرد یکسانی دارند و در مقایسه با الگوریتم بیشترین شباهت (MLP) از دقت بیشتری برخوردار هستند. مطالعات مختلف از جمله باکاری (۲۰۲۴)، جاسیم (۲۰۲۴)، ساویتا (۲۰۲۴) و قدسی و همکاران (۱۳۹۹) هم‌راستا با نتایج پژوهش پیشرو اثربخشی SVM را در دستیابی به نرخ‌های دقت بالا در طبقه‌بندی کاربری و پوشش زمین (LULC) نشان داده‌اند و به این نتیجه رسیده‌اند که عملکرد الگوریتم ماشین بردار پشتیبان (SVM) در طبقه‌بندی کاربری اراضی به‌طور کلی نسبت به سایر الگوریتم‌های محبوب مانند جنگل تصادفی (RF) و حداکثر احتمال (MLC) برتر است (Savitha et al., ۲۰۲۴; Jasim et al., ۲۰۲۴; Baccari et al., ۲۰۲۴).

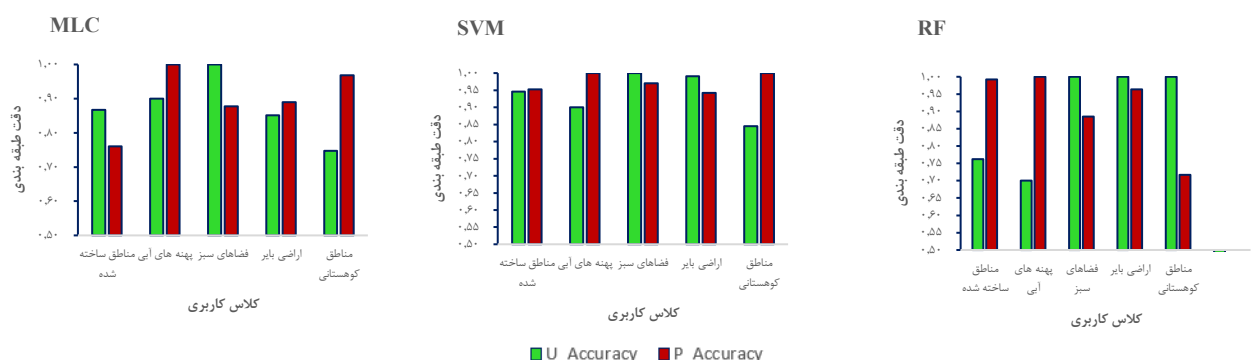
۲۰۲۰, Godsi et al., ۲۰۲۴). در مقابل رازافینیمارو (۲۰۲۲) و ستیانو (۲۰۲۴) بر خلاف پژوهش حاضر، بر دقت بیشتر جنگل تصادفی (RF) نسبت به ماشین بردار پشتیبان (SVM) تأکید کرده‌اند (Setyanto et al., ۲۰۲۲; Razafinimaro et al., ۲۰۲۴).

به طور خلاصه، در حالی که SVM و RF به طور مداوم دقت بالایی را در طبقه‌بندی LULC نشان می‌دهد، انتخاب الگوریتم ممکن است به شرایط خاص و ویژگی‌های داده و محدوده مورد مطالعه بستگی داشته باشد؛ لذا چالش‌های شناسایی کاربری‌ها در منطقه کلان‌شهری مشهد، توسط الگوریتم‌های مورد مطالعه این پژوهش نیز بررسی شده است. در این راستا، کاربری‌های طبقه‌بندی شده با تصاویر گوگل ارث مورد مقایسه قرار گرفته است. در این خصوص، آگاهی و شناخت پژوهشگران نسبت به محدوده مورد مطالعه نیز تأثیر به سزایی داشته است. این مقایسه نشان داد که شباهت رنگ برخی عوارض و مناطق منجر به خطای طبقه‌بندی می‌گردد. به طور مثال برخی از قسمت‌های ساخته شده به عنوان مناطق کوهستانی و برخی از سلول‌های کوهستانی به عنوان مناطق ساخته شده شناسایی شده است. هر دو نوع خطای نوع اول و دوم در این الگوریتم در مورد طبقات مناطق ساخته شده و کوهستانی بیش از دو الگوریتم دیگر است به گونه‌ای که برخی سلول‌هایی که در طبقه مناطق ساخته شده قرار نمی‌گرفتند، به عنوان مناطق ساخته شده شناسایی شدند. این خطا در مورد پهنه‌های آبی نیز صدق می‌کند؛ یعنی سلول‌هایی شناسایی شده‌اند که به اشتباه در طبقه پهنه‌های آبی قرار گرفته‌اند. از طرفی دقت این الگوریتم در شناسایی مناطق کوهستانی نیز کمتر از سایر الگوریتم‌ها است زیرا نتوانسته است مناطق کوهستانی را به دقت شناسایی نماید. بررسی فضایی سلول‌های دارای خطا نشان می‌دهد که برخی از سلول‌هایی که در طبقه مناطق ساخته شده قرار می‌گیرند، به رنگ سوخته هستند و در طیف رنگی مناطق کوهستانی قرار دارند. خطاهای صورت گرفته در خصوص این نوع از سلول‌ها است.

باتوجه به شکل ۵، کمترین خطای نوع اول و دوم در الگوریتم SVM وجود دارد. خطاهای موجود در این الگوریتم نیز بیشتر در شناسایی مناطق کوهستانی است. به گونه‌ای که تعدادی از سلول‌هایی که در مناطق کوهستانی نبودند به عنوان مناطق کوهستانی شناسایی شده‌اند. اما در مقابل خطای نوع دوم در مورد این کاربری در الگوریتم SVM کمتر دیده شده است، یعنی مناطق کوهستانی واقعی، با دقت بسیار بالایی شناسایی شده‌اند.

در الگوریتم MLC به عنوان کم‌دقت‌ترین الگوریتم، علاوه بر مناطق کوهستانی و مناطق ساخته شده، شناسایی فضاهای سبز و اراضی بایر نیز با چالش مواجه بود. به گونه‌ای که برخی از سلول‌های غیر بایر به اشتباه در طبقات بایر قرار گرفتند. علاوه بر این، این الگوریتم در شناسایی فضاهای سبز نیز با خطای بیشتری نسبت به سایر الگوریتم‌ها همراه بوده است.

به طور کلی می‌توان گفت که کاربری بایر با دقت بیشتر و خطای کمتر، و کاربری ساخته شده و مناطق کوهستانی با خطای بیشتری شناسایی شده‌اند.



شکل ۵. خطای نوع اول و دوم در الگوریتم‌های مورد مطالعه

۶- نتیجه گیری

این مطالعه به منظور بررسی دقت طبقه‌بندی پرکاربردترین الگوریتم‌های طبقه‌بندی کاربری اراضی در مشاهدات ماهواره‌ای انجام شد. هدف این پژوهش پیشنهاد بهترین الگوریتم طبقه‌بندی کننده بود.

با بررسی ویژگی‌های مختلف الگوریتم‌های پرکاربرد در طبقه‌بندی کاربری اراضی، سه الگوریتم نظارت شده ماشین بردار پشتیبان، جنگل تصادفی و بیشترین احتمال بر روی داده‌های (OLI) Landsat ۸ برای طبقه‌بندی کاربری اراضی منطقه کلان‌شهری مشهد در سال ۲۰۲۰ انتخاب شدند. ارزیابی صحت طبقه‌بندی کاربری‌های مختلف با استفاده از ضرایب $U_Accuracy$ و $P_Accuracy$ در ماتریس اختلاط و بررسی ضریب تغییرات سهم هر کاربری در الگوریتم‌های مورد مطالعه انجام شد. نتایج بررسی ضرایب $U_Accuracy$ و $P_Accuracy$ نشان می‌دهد که به‌طور کلی صحت طبقه‌بندی طبقات در تمام الگوریتم‌های مورد مطالعه در بازه بین خوب تا عالی قرار می‌گیرد. اما بررسی دقیق‌تر این الگوریتم‌ها نشان می‌دهد که بیش‌ترین چالش شناسایی طبقه برای مناطق ساخته شده، مناطق کوهستانی و فضاهای سبز است و شناسایی اراضی بایر با چالش کمتری مواجه است. همچنین بررسی ضریب تغییرات (CV) سهم طبقات شناسایی شده در الگوریتم‌های مختلف نشان می‌دهد که فضاهای سبز و پهنه‌های آبی از همبستگی کم‌تری در الگوریتم‌های مورد مطالعه برخوردار هستند. ضریب کاپا و تحلیل‌های مبتنی بر ماتریس اختلاط نیز تنوع در دقت هر طبقه‌بندی کننده LULC را نشان می‌دهد. تفاوت در دقت طبقه‌بندی کننده‌های مورد استفاده جزئی است، اما این تغییرات جزئی اهمیت بسیار مهمی در زمینه برنامه‌ریزی LULC دارد. باتوجه به این که این اختلافات جزئی در کاربری‌های حساسی مانند مناطق ساخته شده و فضاهای سبز دیده می‌شود، لذا انتخاب الگوریتمی با بیشترین دقت و کمترین خطا از اهمیت ویژه‌ای برخوردار است. نتایج بررسی ضریب کاپا و تحلیل‌های مبتنی بر ماتریس اختلاط نشان می‌دهد که ماشین بردار پشتیبان بالاترین دقت و کمترین خطا را در بین طبقه‌بندی کننده‌های مورد مطالعه دارد.

باتوجه به این که مطالعات متعدد نشان داد که دقت نقشه‌برداری LULC با زمان و مکان در ارتباط است؛ بنابراین، برای تحقیقات آینده، آنالیز دقت طبقه‌بندی کننده‌ها برای شرایط مورفولوژیک و ژئومورفیک متفاوت پیشنهاد می‌شود.

- Abburu, S., & Golla, S. B. (۲۰۱۵). Satellite image classification methods and techniques: A review. *International journal of computer applications*, 119(۸).
- Aburas, M. M., Ho, Y. M., Ramli, M. F., & Ash'aari, Z. H. (۲۰۱۶). The simulation and prediction of spatio-temporal urban growth trends using cellular automata models: A review. *International Journal of Applied Earth Observation and Geoinformation*, 52, ۳۸۰-۳۸۹. <https://doi.org/10.1016/j.jag.2016.07.007>
- Aljanabi, F., Dedeoğlu, M., & Şeker, C. (۲۰۲۴). Environmental monitoring of Land Use/Land Cover by integrating remote sensing and machine learning algorithms. *Journal of Engineering and Sustainable Development*, 28(۴), ۴۵۵-۴۶۶. <https://doi.org/10.31272/jeasd.28.4.4>.
- Azizi Qalati, S., Rangzen, K., Taghizadeh, A., & Ahmadi, S. (۲۰۱۳). Modeling land use changes using logistic regression method in LCM model (Case study: Kohmera Sorkhi region of Fars province). *Iran Forest and Poplar Research* 22(۴):۵۸۵-۵۹۶.
- BACCARI, N., HAMZA, M. H., SLAMA, T., SEBEL, A., & REBAI, N. (۲۰۲۴). Evaluation of SVM and RF Machine Learning Algorithms in Land Use/Land Cover Change Assessment: Tessa Watershed Case Study (Northwest of Tunisia). *Research Square*. <https://doi.org/10.21203/rs.3.rs-4359112/v1>.
- Bland, J. M., & Altman, D. G. (۱۹۹۵). Multiple significance tests: the Bonferroni method. *Bmj*, 310(۶۹۷۳), ۱۷۰. <https://doi.org/10.1136/bmj.310.6973.170>.
- Chandra, M. A., & Bedi, S. S. (۲۰۲۱). Survey on SVM and their application in image classification. *International Journal of Information Technology*, 13(۵), ۱-۱۱. <https://doi.org/10.1007/s41870-021-0801-1>.
- Dahlke, J., Bogner, K., Mueller, M., Berger, T., Pyka, A., & Ebersberger, B. (۲۰۲۰). Is the juice worth the squeeze? machine learning (ml) in and for agent-based modelling (abm). *arXiv preprint arXiv:2003.11985*. <https://doi.org/10.48550/arXiv.2003.11985>
- de Espindola, G. M., da Costa Carneiro, E. L. N., & Façanha, A. C. (۲۰۱۷). Four decades of urban sprawl and population growth in Teresina, Brazil. *Applied Geography*, 79, ۷۳-۸۳. <https://doi.org/10.1016/j.apgeog.2017.12.018>
- Esmaili, A., Ashjazi, H. (۲۰۱۶). Modeling land use changes through Markov chain and using geographic information systems and remote sensing (case study: Qom province). *Geography and Urban-Regional Studies* 9(۳۱):۱۷۲-۱۵۳. ۱۰,۲۲۱۱/GAIJ.۲۰۱۹,۴۷۱.
- Fleiss, J. L. (۱۹۷۱). Measuring nominal scale agreement among many raters. *Psychological bulletin*, 76(۵), ۳۷۸. <https://doi.org/10.1037/h0031619>.
- Gan, H., Sang, N., Huang, R., Tong, X., & Dan, Z. (۲۰۱۳). Using clustering analysis to improve semi-supervised classification. *Neurocomputing*, 101, ۲۹۰-۲۹۸. <https://doi.org/10.1016/j.neucom.2012.08.020>
- Ganesh, B., Vincent, S., Pathan, S., & Benitez, S. R. G. (۲۰۲۳, October). Utilizing LANDSAT data and the Maximum Likelihood Classifier for Analysing Land Use Patterns in Shimoga, Karnataka. In *Journal of Physics: Conference Series* (Vol. ۲۵۷۱, No. ۱, p. ۰۱۲۰۰۱). IOP Publishing. ۱۰.۱۰۸۸/۱۷۴۲-۶۵۹۶/۲۵۷۱/۱/۰۱۲۰۰۱.
- Qudsi Khairkhash Zarkash, M., & Khirzmecheshme, Baqer. (۲۰۲۱). Comparison of accuracy of support vector machine and random forest methods in preparing land use map and crops, using Sentinel-۲ multi-temporal images. *Iranian Remote Sensing and GIS Journal* 12(۴), ۷۳-۹۲. ۱۰,۵۲۵۴۷/GISJ.۱۲,۴,۷۳.

Gholamali Fard, M., Jorabian Shoshtri, SH., Kohnouj, H., & Mirzaei, M. (۲۰۱۱). Modelling land use changes of the coasts of Mazandaran province using LCM in GIS environment. *Ecology* 38(۶۴): ۱۰۹-۱۲۴.

Glielmo, A., Husic, B. E., Rodriguez, A., Clementi, C., Noé, F., & Laio, A. (۲۰۲۱). Unsupervised learning methods for molecular simulation data. *Chemical Reviews*, 121(۱۶), ۹۷۲۲-۹۷۵۸. <https://doi.org/10.1021/acs.chemrev.1c01190>.

Hosseini, M., Karimi, M., Saadi Masgari, M., & Heydari, M. (۲۰۱۶). Design and implementation of an integrated system of urban land use change modeling. *Geographical Sciences* 16(۴۰): ۶۹-۹۱.

Jasim, B. S., Jasim, O. Z., & AL-Hameedawi, A. N. (۲۰۲۴). Evaluating land use land cover classification based on machine learning algorithms. *Engineering and Technology Journal*, 42(۵), ۵۵۷-۵۶۸. <http://doi.org/10.30684/etj.2024.144080.1638>

Jiao, S. R., Song, J., & Liu, B. (۲۰۲۰, November). A review of decision tree classification algorithms for continuous variables. In *Journal of Physics: Conference Series* (Vol. ۱۶۵۱, No. ۱, p. ۰۱۲۰۸۳). IOP Publishing ۱۰.۱۰۸۸/۱۷۴۲-۶۵۹۶/۱۶۵۱/۱.۱۲۰۸۳.

Landis, J. R., & Koch, G. G. (۱۹۷۷). An application of hierarchical kappa-type statistics in the assessment of majority agreement among multiple observers. *Biometrics*, ۳۶۳-۳۷۴. <https://doi.org/10.23۰۷/۲۵۲۹۷۸۶>

Louppe, G. (۲۰۱۴). Understanding random forests. *Cornell University Library*, 10.

Mahmoudzadeh, H., Derakhshani, K., & Momeni, S. (۲۰۱۹). Modeling the effects of marginalization on the changes of Urmia city and predicting the physical development of the city using satellite images until the horizon of ۱۴۱۰. *Human Geography Research* 51(۴): ۹۰-۸۷۱. ۱۰.۲۲۰۵۹/JHGR.۲۰۱۸.۲۳.۷۸۸,۱۰.۷۴۳۴

Marudi, M., Ben-Gal, I., & Singer, G. (۲۰۲۴). A decision tree-based method for ordinal classification problems. *IJSE Transactions*, 56(۹), ۹۶۰-۹۷۴. <https://doi.org/10.1080/24720804.2022.2081740>

Mashhad Municipality. (۲۰۱۶). *Statistics of Mashhad city*. Vice President of Planning and Human Capital Development of Mashhad Municipality with the supervision of statistics management, performance analysis and evaluation.

Mendoza, M. E., Granados, E. L., Geneletti, D., Pérez-Salicrup, D. R., & Salinas, V. (۲۰۱۱). Analysing land cover and land use change processes at watershed level: a multitemporal study in the Lake Cuitzeo Watershed, Mexico (۱۹۷۵-۲۰۰۳). *Applied Geography*, 31(۱), ۲۳۷-۲۵۰. <https://doi.org/10.1016/j.apgeog.2010.05.010>

Mirakhorlou, M., & Rahimzadegan, S. Land use change modeling using Markov-cellular automata model and multi-criteria decision making in Talar watershed. (۲۰۱۸). *Scientific Journal of Mapping Sciences and Techniques* 8(۱): ۸۵-۹۹

Mohammadi, J., & Mahmoud, A. (۲۰۱۲). Spatial analysis and urban land use planning of Dogonbadan (Gachsaran). *Geographical Research* 27(۲): ۱۹-۳۶.

Moradi, Mohammad, & Kasani, E. (۲۰۲۴). Optimal Classification Using Decision Tree. *Mathematical Researches* 10(۲): ۵۱-۶۵.

Naswin, A., & Wibowo, A. P. (۲۰۲۳). Performance Analysis of the Decision Tree Classification Algorithm on the Pneumonia Dataset. *International Journal of Artificial Intelligence in Medical Issues*, 1(۱), ۱-۹.

Nti, I. K., Nyarko-Boateng, O., Adekoya, A. F., Weyori, B. A., & Adjei, H. P. (۲۰۲۳). Predicting diabetes using Cohen's Kappa blending ensemble learning. *International Journal of Electronic Healthcare*, 13(۱), ۵۷-۷۰. <https://doi.org/10.1002/IJEH.2023.1286.0>.

Piryonesi, S. M. (۲۰۱۹). *The application of data analytics to asset management: Deterioration and climate change adaptation in Ontario roads*. University of Toronto (Canada).

Pourkhabaz, H., Mohammadyari, F., Aqdar, H., & Tavakoli, M. An experimental approach in modeling land use changes in Behbahan city using multi-temporal satellite images *Town & Country Planning* 7(۲):2008-7047. ۱۰,۲۲۰۵۹/JTCP.۲۰۱۵,۵۷۱۸۷.

Qiang, W., & Zhongli, Z. (۲۰۱۱, August). Reinforcement learning model, algorithms and its application. In *2011 International Conference on Mechatronic Science, Electric Engineering and Computer (MEC)* (pp. ۱۱۴۳-۱۱۴۶). IEEE ۱۰,۱۱۹/MEC.۲۰۱۱,۶۰۲۵۶۶۹.

Razafinimaro, A., Hajalalaina, A. R., Rakotonirainy, H., & Zafimarina, R. (۲۰۲۲). Land cover classification based optical satellite images using machine learning algorithms. *International Journal of Advances in Intelligent Informatics*, 8(۳), ۳۶۲-۳۸۰. <https://doi.org/۱۰,۲۶۵۵۵/ijain.v۸i۳.۸۰۳>

Ritsema van Eck, J., & Koomen, E. (۲۰۰۸). Characterising urban concentration and land-use diversity in simulations of future land use. *The Annals of Regional Science*, 42, ۱۲۳-۱۴۰. <https://doi.org/۱۰,۱۰۰۷/s۰۰۱۶۸-۰۰۷-۰۱۴۱-۷>

Sakai, T., Plessis, M. C., Niu, G., & Sugiyama, M. (۲۰۱۷, July). Semi-supervised classification based on classification from positive and unlabeled data. In *International conference on machine learning* (pp. ۲۹۹۸-۳۰۰۶). PMLR.

Santarsiero, V., Nolè, G., Lanorte, A., Tucci, B., Cillis, G., & Murgante, B. (۲۰۲۲). Remote sensing and spatial analysis for land-take assessment in Basilicata Region (Southern Italy). *Remote Sensing*, 14(۷), ۱۶۹۲. <https://doi.org/۱۰,۳۳۹۰/rs۱۴۰۷۱۶۹۲>

Savitha, C., & Reshma, T. (۲۰۲۲). Performance Evaluation of Support Vector Machine and Random Forest Techniques for Land Use-Land Cover Classification—A Case Study on a Mili Scale Agricultural Watershed, Tadepalligudem, India. In *International Virtual Conference on Developments and Applications of Geomatics:۳۷۹-۳۹۲*. Singapore: Springer Nature Singapore.

Setyanto, A., & Sunyoto, A. (۲۰۲۴). Comparative Study SVM and Random Forest Algorithms for the Classification of Terrestrial Visual Rock Types. In *IOP Conference Series: Earth and Environmental Science* 1357(۱), ۱۰,۱۰۸۸/۱۷۵۵-۱۳۱۵/۱۳۵۷/۱,۰۱۲,۳۶.

Shi, S., Zhang, X., & Fan, W. (۲۰۱۹). Explaining the predictions of any image classifier via decision trees. *arXiv preprint arXiv:1911.01058*. <https://doi.org/۱۰,۴۸۵۵/arXiv.۱۹۱۱,۰۱۰۵۸>

Sowmya, D. R., Shenoy, P. D., & Venugopal, K. R. (۲۰۱۷). Remote sensing satellite image processing techniques for image classification: a comprehensive survey. *International Journal of Computer Applications*, 161(۱۱), ۲۴-۳۷.

Talukdar, S., Singha, P., Mahato, S., Pal, S., Liou, Y. A., & Rahman, A. (۲۰۲۰). Land-use land-cover classification by machine learning classifiers for satellite observations—A review. *Remote sensing*, 12(۷), ۱۱۳۵. <https://doi.org/۱۰,۳۳۹۰/rs۱۲۰۷۱۱۳۵>

Yimer, S. M., Bouanani, A., Kumar, N., Tischbein, B., & Borgemeister, C. (۲۰۲۴). Comparison of different machine-learning algorithms for land use land cover mapping in a heterogenous landscape over the Eastern Nile river basin, Ethiopia. *Advances in Space Research*, 74(۵), ۲۱۸۰-۲۱۹۹. <https://doi.org/۱۰,۱۰۱۶/j.asr.۲۰۲۴,۰۶,۰۱>

Yin, M., Wortman Vaughan, J., & Wallach, H. (۲۰۱۹, May). Understanding the effect of accuracy on trust in machine learning models. In *Proceedings of the 2019 chi conference on human factors in computing systems* (pp. ۱-۱۲). <https://doi.org/۱۰,۱۱۴۵/۳۲۹.۶۰۵,۳۳.۰۵۰۹>