

Original Article

Iranian
Journal of Remote Sensing
GIS

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 Rahnama^{*ID}, Mohammad Ajza Shokuhi, Hossein Vahidi

Affiliation

Dep. of Geography, Faculty of Literature and Human Sciences, Ferdowsi University of Mashhad, Mashhad, Iran

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 0.93, 0.88, and 0.80, 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.

Keywords: Remote sensing, Land use classification, Support Vector Machine, Random Forest, Maximum Likelihood.

Citation:

Baghban, S., Rahnama, M. R., Ajza Shokuhi, M., Vahidi, H., Performance Analysis of Support Vector Machine, Random Forest, and Maximum Likelihood Algorithms in Land Use Classification of the Metropolitan Area of Mashhad, Iran J Remote Sens GIS. 17(4): 111-132.

* Corresponding Author: rahnama@um.ac.ir
DOI: <https://doi.org/10.48308/gisj.2024.236864.1230>

Received: 2024.09.10

Accepted: 2024.11.10



Copyright: © 2026 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license <https://creativecommons.org/licenses/by/4.0/>.



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

ساجده باغبان، محمدرحیم رهنما^{*}، محمد اجزاء شکوهی، حسین وحیدی

چکیده

سابقه و هدف: از آنجاکه ارزش و امکان استفاده از هر نقشه تولیدشده براساس تصاویر ماهواره‌ای با توجه به میزان صحت آن مشخص می‌شود، ارزیابی صحت روش طبقه‌بندی تصاویر ماهواره‌ای دارای اهمیت چشمگیری است. از این‌رو این پژوهش با هدف تحلیل کارایی الگوریتم‌های ماشین بردار پشتیبان (SVM)، جنگل تصادفی (RF) و حداکثر احتمال (MLC)، در شناسایی کاربری و پوشش اراضی منطقه کلان‌شهری مشهد انجام شده است. تا به امروز الگوریتم‌های بسیار زیادی، به‌منظور طبقه‌بندی تصاویر ماهواره‌ای، توسعه یافته‌اند که عملکرد آن‌ها، در شرایط گوناگون، متفاوت است. به‌همین دلیل در این پژوهش، ابتدا با مروری بر پژوهش‌های پیشین، پراکندترین الگوریتم‌ها شناسایی شده و سپس، با سنجش ویژگی‌های انواع طبقه‌بندی‌کننده‌ها، سه الگوریتم ماشین بردار پشتیبان و جنگل تصادفی و حداکثر احتمال انتخاب شده است. با توجه به اینکه مطالعات متعدد نشان داده است دقت نقشه‌برداری LULC تحت تأثیر زمان و مکان قرار دارد و هر یک از پژوهش‌های انجام‌شده نیز بر دقت الگوریتم‌های متفاوتی تأکید کرده‌اند، نتایج آن‌ها درمورد شرایط جغرافیایی ایران تعمیم‌پذیر نیست. از طرفی، در شرایط ژئومورفولوژیک ایران، پژوهش‌های کافی به‌منظور سنجش دقت الگوریتم‌های طبقه‌بندی انجام نشده و اغلب مطالعات صحت‌سنجی الگوریتم‌ها در نمونه‌های موردی خارج از ایران انجام شده است. از این‌رو با توجه به تفاوت نتایج الگوریتم‌ها در شرایط گوناگون، بررسی دقت و عملکرد الگوریتم‌ها با تمرکز بر منطقه وسیع و متنوع کلان‌شهری مشهد می‌تواند نتایج بدیع و جالب‌توجهی به‌همراه داشته باشد.

سمت

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

مواد و روش‌ها: روش تحقیق حاضر، ازمنظر هدف، کاربردی و ازمنظر ماهیت، توصیفی-تحلیلی است. گردآوری اطلاعات در این پژوهش به‌روش اسنادی-کتابخانه‌ای انجام شده است. در این مطالعه، تصویر سنجده OLI در ماهواره لندست-۸ تهیه شده است. طبقه‌بندی تصاویر ماهواره‌ای در دو مرحله پیش‌پردازش 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 مشخص شد؛ این کاربری‌ها شامل پنج دسته و بدین‌قرار است: (۱) مناطق ساخته‌شده؛ (۲) اراضی بایر؛ (۳) مناطق کوهستانی؛ (۴) فضاهاى سبز؛ (۵) پهنه‌های آبی.

نتایج و بحث: نتایج حاصل ارزیابی انحراف معیار (SD) و ضریب تغییرات (CV) درصد سهم مساحت در یک کلاس LULC با استفاده از الگوریتم‌های گوناگون نشان می‌دهد که اراضی بایر با دقت بیشتر و پهنه‌های آبی و فضاهاى سبز با دقت کمتری طبقه‌بندی شده‌اند. نتایج بررسی ضرایب P_Accuracy و U_Accuracy نشان می‌دهد که به‌طور کلی، صحت طبقه‌بندی دسته‌ها در تمامی الگوریتم‌های مورد مطالعه، در بازه خوب تا عالی قرار می‌گیرد. اما بررسی دقیق‌تر این الگوریتم‌ها نشان می‌دهد که بیشترین چالش شناسایی طبقه‌ها درمورد مناطق ساخته‌شده، مناطق کوهستانی و فضاهاى سبز وجود دارد 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 چشمگیری در زمینه برنامه‌ریزی LULC دارد. با توجه به اینکه این اختلافات جزئی در کاربری‌های حساسی، مانند مناطق ساخته‌شده و فضاهاى سبز دیده می‌شود، انتخاب الگوریتمی دارای بیشترین دقت و کمترین خطا اهمیت ویژه‌ای دارد.

استناد:

باغبان، س.، رهنما، م.ر.، اجزاء شکوهی، م.، وحیدی، ح.، تحلیل کارایی الگوریتم‌های ماشین بردار پشتیبان، جنگل تصادفی و حداکثر احتمال در شناسایی کاربری اراضی منطقه کلان‌شهری مشهد، نشریه سنجش از دور و GIS ایران، سال ۱۷، شماره ۴، زمستان ۱۴۰۴، ۱۱۱-۱۳۲.

نتیجه‌گیری: نتایج بررسی ضریب کاپا و تحلیل‌های مبتنی بر ماتریس اختلاط نشان می‌دهد که رویکرد SVM دقت کلی بیشتر و ضریب کاپای بالاتری از روش‌های RF و MLC دارد؛ به‌گونه‌ای که الگوریتم‌های SVM، RF و MLC به‌ترتیب، دقت کلی معادل ۰/۹۳، ۰/۸۸ و ۰/۸۰ را به دست آورده‌اند. بنابراین ماشین بردار پشتیبان بیشترین دقت و کمترین خطا را در بین طبقه‌بندی‌کننده‌های مورد مطالعه دارد. براین‌اساس که مطالعات متعدد گویای ارتباط میان دقت نقشه‌برداری LULC با زمان و مکان است، درمورد تحقیقات آینده، تحلیل دقت طبقه‌بندی‌کننده‌ها برای شرایط مورفوکلیماتیک و ژئومورفیک متفاوت پیشنهاد می‌شود.

واژه‌های کلیدی: سنجش از دور، طبقه‌بندی کاربری اراضی، ماشین بردار پشتیبان، جنگل تصادفی، بیشترین احتمال، مشهد.



۱- مقدمه

کاربری اراضی، به منزله یکی از اجزای سیستم شهری، فرایندهای دینامیک مکانی- زمانی را شامل می‌شود که با توجه به ارزش زمین، در کنترل درآوردن آن دارای اهمیت است؛ بنابراین آگاهی از تغییرات و تحولات کاربری اراضی و عوامل مؤثر در آن، طی دوره‌ای زمانی، برای برنامه‌ریزان و مدیران بسیار مهم است و چه بسا هدف مستقیم برنامه‌ریزی فضایی در نظر گرفته شود. به همین دلیل استفاده از روش‌های آشکارسازی تغییرات، برای مشخص کردن روند تغییرات با گذشت زمان، ضروری به نظر می‌رسد (Mohammadi & Akbari, 2018; Hosseini et al., 2016; Aburas et al., 2016; Ritsema van Eck & Koomen, 2008) و در نتیجه، تهیه لایه پوشش و کاربری اراضی در گذشته و شرایط فعلی، و روند تغییرات آن می‌تواند شناخت دقیق از چندوچون تغییرات منطقه به دست دهد. در این راستا تکنولوژی‌های نوین، مانند سنجش از دور و سیستم اطلاعات جغرافیایی، ابزارهای مناسبی برای اندازه‌گیری وسعت و الگوی تغییرات، طی زمان فراهم کرده است. نظارت بر رشد و توسعه شهری به کاربردی مهم در زمینه داده‌های سنجش از دور و سیستم اطلاعات مکانی تبدیل شده است؛ به طوری که روند تعیین تغییرات در کاربری زمین، براساس چند دوره داده‌های سنجش از دور، امکان‌پذیر است. استفاده از تکنیک‌های سنجش از دور می‌تواند سطح و روند تغییرات پوشش اراضی را با هزینه کمتر و سرعت بیشتر، امکان‌پذیر کند. تصاویر ماهواره‌ای، به دلیل قدرت تفکیک مکانی بالا، تکرارپذیر بودن و امکان پردازش‌های رایانه‌ای، از ابزارهای مهم در تهیه نقشه‌های کاربری اراضی و نیز مطالعه تغییرات زمانی در مناطق گوناگون شمرده می‌شود؛ در این روش، می‌توان کاربری اراضی را برحسب تصاویر چندزمانه تفکیک کرد و از روش مقایسه پس از طبقه‌بندی، برای پایش تغییرات، بهره برد تا به شناخت مناسبی درباره چگونگی تغییرات کاربری اراضی دست یافت (Esmaili & Ashjaei, 2019; De, 2011; Espindola et al., 2017; Mendoza et al., 2011).

اما با توجه به اینکه ارزش و امکان استفاده از هر نقشه تولیدشده براساس تصاویر ماهواره ای به میزان صحت آن بازمی‌گردد، ارزیابی صحت روش طبقه‌بندی تصاویر ماهواره‌ای اهمیت بسیاری دارد. صحت نقشه‌ها طبق فاکتورهایی همچون صحت داده‌های اولیه و روش طبقه‌بندی تصویر ماهواره‌ای مشخص می‌شود. اصولاً طبقه‌بندی پیکسل‌ها در دسته‌های مشخص، براساس ارزش‌های بازتابشی ثبت‌شده در فضای هر تصویر ماهواره‌ای است. در عمل، طبقه‌بندی هریک از مقادیر روشنایی‌ها به کلاس‌های پوشش اراضی، زمین‌شناسی، کاربری اراضی و دیگر عوارض سطح زمین منتسب می‌شود (Abburu & Golla, 2015; Sowmya et al., 2017).

الگوریتم‌های بسیار زیادی، به منظور طبقه‌بندی تصاویر ماهواره‌ای، تا به امروز توسعه یافته‌اند که عملکرد آن‌ها، در شرایط گوناگون، متفاوت است؛ از این رو در این پژوهش، ابتدا با مروری بر پژوهش‌های پیشین، پرکاربردترین الگوریتم‌ها شناسایی شد. 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 هریک از آن‌ها ارزیابی شد. به طور کلی تنوع در رویکردها، کارایی در مسائل گوناگون، سازگاری با داده‌های پیچیده، امکان 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 داشته است.

پژوهش‌های تالوکدر^۷ و همکاران (۲۰۲۰)، سانتارسیهرو^۸ و همکاران (۲۰۲۲) و چندرا و بدی^۹

1. Support Vector Machine (SVM)
2. Random Forest (RF)
3. Decision Tree (DT)
4. Maximum Likelihood (ML)
5. K-Nearest Neighbors (KNN)
6. Artificial Neural Networks ANN
7. Talukdar
8. Santarsiero
9. Chandra & Bedi

محبوب‌ترین الگوریتم‌ها بررسی و سپس سعی شده است، با سنجش ویژگی‌های انواع طبقه‌بندی‌کننده‌ها، دلایل انتخاب آن‌ها به‌طور شفاف بیان شود. همچنین با توجه به اینکه مطالعات متعدد نشان داده است دقت نقشه‌برداری LULC با زمان و مکان در ارتباط است و هریک از پژوهش‌های انجام‌شده نیز بر دقت الگوریتم‌های متفاوتی تأکید داشته، بنابراین نتایج آن‌ها در مورد شرایط جغرافیایی ایران تعمیم‌پذیر نیست. از طرفی، در شرایط ژئومورفولوژیک ایران، پژوهش‌های کافی برای سنجش دقت الگوریتم‌های طبقه‌بندی انجام نشده و در بسیاری از مطالعات، صحت‌سنجی الگوریتم‌ها روی نمونه‌های موردی خارج از ایران انجام شده است. این مسئله گویای تفاوت‌های فرهنگی، جغرافیایی و زیست‌محیطی است که می‌تواند در نحوه عملکرد الگوریتم‌ها تأثیرگذار باشد. بدین ترتیب و با وجود این تفاوت‌ها، بررسی دقت و عملکرد الگوریتم‌های گوناگون در شرایط خاص و بی‌نظیر کلان‌شهری مشهد می‌تواند، در حوزه طبقه‌بندی کاربری اراضی، به نتایج جدید و بدیع منجر شود.

علاوه‌براین در پژوهش‌های قبلی، به چالش‌های خاص در زمینه طبقه‌بندی کاربری‌های گوناگون، به‌طور جامع پرداخته نشده است. این چالش‌ها می‌تواند شامل تنوع کاربری‌ها، پیچیدگی الگوهای جغرافیایی و تغییرات مستمر در استفاده از اراضی باشد. تحقیق پیش رو، با شناسایی و تحلیل این چالش‌ها در الگوریتم‌های گوناگون، سعی دارد به پژوهشگران و تصمیم‌گیرندگان کمک کند که در مطالعات آتی خود، براساس کلاس‌های طبقه‌بندی و ویژگی‌های منطقه‌ای، الگوریتم‌های مناسب‌تری را برگزینند.

(۲۰۲۱) نشان می‌دهد که روش ماشین بردار پشتیبان دارای دقت معتناهی است. همچنین الگوریتم جنگل تصادفی، در پژوهش‌های مرادی و کسانی^۱ (۲۰۲۴)، شی^۲ و همکاران (۲۰۱۹)، نسوین و ویوو^۳ (۲۰۲۳)، ژیانو^۴ و همکاران (۲۰۲۰) تأیید شده است. بررسی پژوهش‌های انجام‌شده در زمینه طبقه‌بندی کاربری اراضی نشان می‌دهد که الگوریتم حداکثر احتمال بیشترین فراوانی را در میان الگوریتم‌های مورد استفاده، به‌خصوص در پژوهش‌های داخلی، دارد (Rahnema, 2021; Mahmoudzadeh et al., 2019; Pourkhabbaz et al., 2015; Gholamali Fard et al., 2011; Azizi Qalati et al., 2013; MirAkhrolou & Rahimzadegan, 2018).

در زمینه بررسی عملکرد الگوریتم‌های گوناگون طبقه‌بندی، پژوهش‌های متفاوتی انجام شده که نتایج آن‌ها با هم متفاوت است. در این میان، می‌توان به پژوهش جهانبخشی و اختصاصی^۵ (۲۰۱۹) اشاره کرد که در آن الگوریتم‌های جنگل تصادفی، ماشین بردار پشتیبان و بیشترین شباهت در طبقه‌بندی انواع کاربری کشاورزی مقایسه شده است. در این پژوهش، تصاویر ماهواره لندست-۸ به کار رفته است. در پژوهش قدسی^۶ و همکاران (۲۰۲۱)، به‌منظور طبقه‌بندی تصاویر ماهواره سنتینل برای طبقه‌بندی کاربری اراضی کشاورزی، روش‌های جنگل تصادفی و ماشین بردار پشتیبان استفاده شده است. تالوکدر و همکاران (۲۰۲۰) شش الگوریتم از الگوریتم‌های زیرمجموعه یادگیری ماشین را بررسی کردند. الجنابی^۷ و همکاران (۲۰۲۴) نیز، مانند همین پژوهش، الگوریتم‌های جنگل تصادفی، ماشین بردار پشتیبان و بیشترین احتمال را مورد مطالعه داده‌اند.

یکی از مهم‌ترین معایب پژوهش‌هایی که دقت الگوریتم‌ها را با یکدیگر مقایسه کرده‌اند آن است که در انتخاب الگوریتم‌های مورد مقایسه، هیچ معیاری وجود ندارد. درواقع، با توجه به تعدد الگوریتم‌های موجود، مشخص نیست که پژوهشگر چرا چند الگوریتم خاص را برای مقایسه انتخاب کرده است؛ ازاین‌رو در این پژوهش، به‌منظور رفع نواقص پژوهش‌های پیشین، ابتدا

1. Moradi & Kasani

2. Shi

3. Naswin & Wibowo

4. Jiao

5. Jahanbakhshi & Ekhtesasi

6. Ghodsi

7. Aljanabi

انتخاب کلان‌شهر مشهد در جایگاه منطقه مورد مطالعه نیز، به دلیل تحولات گسترده‌ای که طی چند دهه اخیر در آن رخ داده، از اهمیت ویژه‌ای برخوردار است. مشهد، یکی از مهم‌ترین قطب‌های توسعه در شرق کشور، شاهد تغییرات اجتماعی، اقتصادی و زیست‌محیطی شایان توجهی بوده است که می‌تواند در طبقه‌بندی کاربری اراضی تأثیرگذار باشد.

تنوع جغرافیایی و پویایی‌های این منطقه از عوامل اصلی به چالش کشیدن طبقه‌بندی کاربری‌های گوناگون است؛ به گونه‌ای که با اعمال موفق یک روش در چنین منطقه‌ای، می‌تواند امکان استفاده از آن را در دیگر محدوده‌های مشابه تضمین کند. در نهایت، تنوع الگوریتم‌های موجود برای طبقه‌بندی تصاویر ماهواره‌ای و نیاز به بررسی کارایی آن‌ها، در شرایط جغرافیایی خاص ایران، در نظر گرفته شد و 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 شد. این پژوهش در پی آن بوده است که به‌منظور بهبود روش‌های طبقه‌بندی و ارتقای مدیریت و برنامه‌ریزی اراضی، نتایج کاربردی و مورد اعتمادی عرضه کند.

۲- مواد و روش‌ها

۲-۱- محدوده مورد مطالعه

کلان‌شهر مشهد در $59^{\circ}30'$ تا $60^{\circ}35'$ طول شرقی و $36^{\circ}59'$ تا $35^{\circ}42'$ عرض شمالی واقع شده است و با مساحتی در حدود هجده هزار کیلومتر مربع و بیش از سه میلیون نفر جمعیت (براساس سرشماری سال ۱۳۹۵)، دومین کلان‌شهر ایران به شمار می‌آید (Mashhad Municipality, 2016). پویایی این منطقه کلان‌شهری عامل مهمی محسوب می‌شود که در انتخاب این منطقه، به‌منزله محدوده مورد مطالعه، مؤثر بوده است. عوامل طبیعی و انسانی بسیاری در پویایی محدوده تأثیرگذارند. کمربند سبز مشهد، با ۲۷ کیلومتر طول، در جنوب مشهد قرار دارد که بخشی از آن از

میان کوه‌های جنوب شهر می‌گذرد. هرساله در بخش‌هایی از دامنه‌های این کوه‌ها ساخت‌وساز انجام می‌شود و از بین می‌روند. شهرستان‌های طرقله و شاندیز، با قرار گرفتن در منطقه کلان‌شهری دومین کلان‌شهر ایران، به‌صورت عاملی انسانی در پویایی غرب و جنوب غرب منطقه نقش داشته‌اند. دامنه‌های شمالی رشته‌کوه بینالود از جنوب این منطقه می‌گذرد و در توسعه شهرستان‌های طرقله و شاندیز، به‌منزله مناطق گردشگری، نقش بسزایی دارد؛ از این‌رو زمین‌های کشاورزی و باغ‌های دره‌های این رشته‌کوه در معرض خطر قرار دارند. توسعه نقش گردشگری در این دو شهرستان موجب گسترش و افزایش کاربری‌های تفریحی، گردشگری و تجاری شده و زمین‌های کشاورزی و کوهستانی^۱ منطقه را تحت تأثیر قرار داده است. شدیدترین تغییرات، در کاربری زمین، در شمال شهر مشهد رخ می‌دهد. جریان رودخانه کشف و سرشاخه‌های آن در شمال شهر باعث توسعه زمین‌های کشاورزی در این بخش شده است. با توجه به وجود کوه‌های هزارمسجد در شمال و قرار گرفتن شهرستان‌های طرقله و شاندیز در غرب شهر و به‌دلیل وجود زمین‌های کشاورزی در شمال، تمرکز روستاها در شمال شهر را می‌توان مشاهده کرد. به‌علاوه، با گسترش غیررسمی سکونتگاه‌های انسانی در شمال، شهر در همین راستا توسعه یافته و زمین‌های کشاورزی را تحت تأثیر قرار داده است. پویایی کاربری اراضی در این بخش به‌گونه‌ای است که از طرفی، سکونتگاه‌های انسانی در دل اراضی بایر^۲ و کشاورزی توسعه می‌یابد و از دیگر سو، با افزایش خشکسالی و فقدان مدیریت صحیح منابع طبیعی، اراضی کشاورزی بسیاری از بین می‌روند و به اراضی بایر تبدیل می‌گردند.

تنوع جغرافیایی عامل دیگری است که در انتخاب این منطقه، به‌منزله محدوده مورد مطالعه، نقش داشته است. فضاها سبز^۳ محدوده دارای تنوع بسیاری است؛ به‌گونه‌ای که بخش‌های گوناگون این فضاها پوشش‌های

1. Rocks
2. Barren Lands
3. Green Spaces

۱-۲-۲- جمع‌آوری داده‌ها

در این مطالعه، تصویر سنجنده^۲ OLI در ماهواره^۱ لندست- ۸ (تاریخ ۲۰۲۰/۰۷/۰۸) به‌طور خاص، به دلیل نبود ابر و غبار، از پایگاه داده^۳ USGS Earth Explorer تهیه شده است. تاریخ انتخابی مصادف است با اوایل تابستان که به‌منظور افزایش دقت در شناسایی تفاوت‌ها و تمایز بین مناطق دارای پوشش گیاهی و مناطق اطراف آن، در نظر گرفته شده است. در این زمان، اراضی دارای پوشش گیاهی علاوه‌بر اینکه بدون ابرند، در متراکم‌ترین حالت خود نیز قرار دارند؛ به‌کمک این ویژگی، امکان تحلیل دقیق‌تری از کاربری‌های اراضی فراهم می‌شود و کیفیت تصویر نیز، به‌علت خلوص آن، افزایش می‌یابد.

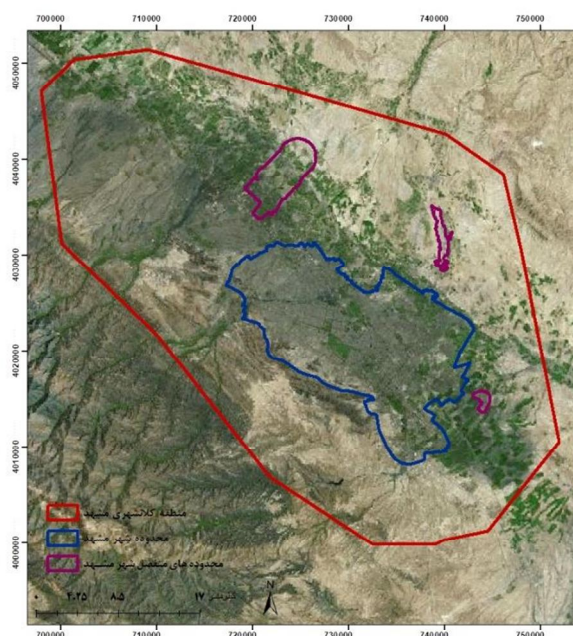
تصاویر ماهواره^۱ لندست، به‌دلیل دوره^۴ تاریخی طولانی‌تر در قیاس با ماهواره^۵ سنتینل، انتخاب شده است. یکی از اهداف کلی طبقه‌بندی کاربری اراضی پایش تغییرات در طول زمان است و تحلیل روند تغییرات کاربری‌ها با تصاویر این ماهواره بهتر انجام می‌شود.

متفاوتی، ازجمله کشاورزی، جنگلی، آیش، متراکم و کم‌تراکم، دارند. اراضی بایر نیز دارای جنس‌های متفاوتی، ازجمله خاک و ماسه‌اند. پهنه‌های آبی^۱ به‌صورت پراکنده در این منطقه مشاهده می‌شود. علاوه‌براین، مناطق صخره‌ای در جنوب محدوده با دره‌های سبز و اراضی بایر آمیخته شده‌اند.

طبقه‌بندی، در محدوده‌ای ایستا و نسبتاً پایدار، قاعداً آسان‌تر از منظره‌هایی است که مانند منطقه^۲ کلان‌شهری مشهد دارای چنین وسعت، تنوع و پویایی چشمگیری است؛ بنابراین عملکرد مطلوب هر الگوریتم، در طبقه‌بندی چنین منطقه‌ای، امکان استفاده و تعمیم آن را اثبات می‌کند.

۲-۲- روش تحقیق

شیوه^۳ این تحقیق، ازمنظر هدف، کاربردی و ازمنظر ماهیت، توصیفی-تحلیلی است و گردآوری اطلاعات در این پژوهش به‌روش اسنادی-کتابخانه‌ای انجام شده است. شکل ۲ نمودار جریان فرایند تحقیق را نشان می‌دهد.



شکل ۱. محدوده مورد مطالعه

1. Water Bodies
2. Operational Land Imager
3. earthexplorer.usgs.gov

ترکیب رنگی کاذب فروسرخ نزدیک^۱، باند ۴ (فروسرخ نزدیک) به صورت کانال قرمز، باند ۳ (قرمز) به صورت کانال سبز و باند ۲ (سبز) به صورت کانال آبی نمایش داده شده است. سپس به این دلیل که وسعت تصاویر ماهواره‌ای بیشتر از منطقه کلان‌شهری مشهد بود، منطقه مورد مطالعه برش خورده و پس از آن، پردازش شده است.

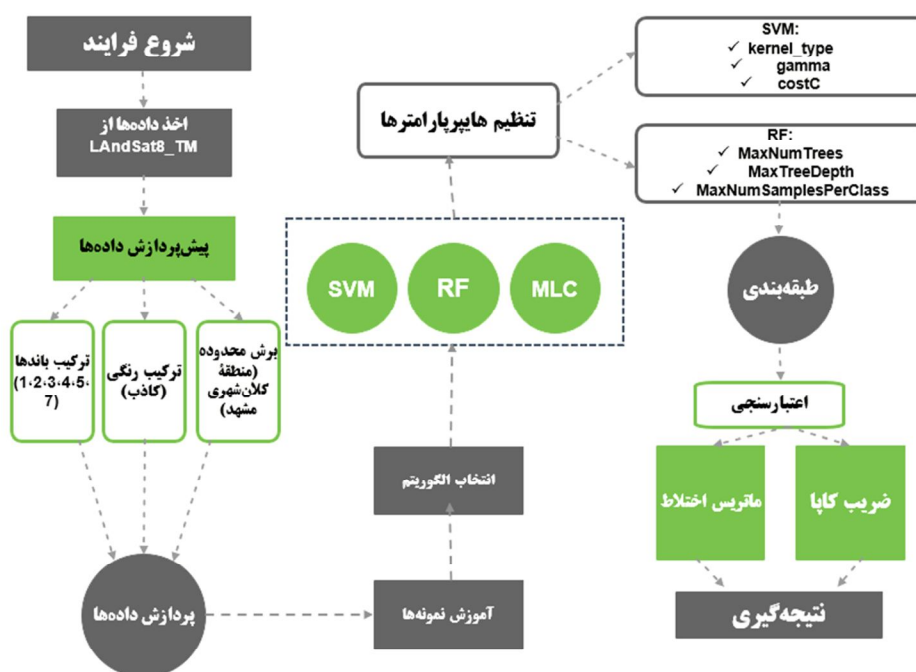
۲-۲-۳- پردازش تصاویر ماهواره‌ای

طبقه‌بندی تصاویر ماهواره‌ای چه‌بسا اصلی‌ترین مرحله پردازش به شمار می‌رود و از این راه، تبدیل فضای تصویر (بازتابش‌های ثبت‌شده در باندهای گوناگون) به فضای واقعی (نقشه‌های پوشش زمین و کاربری اراضی) امکان‌پذیر می‌شود. به جداسازی مجموعه‌های طیفی مشابه و تقسیم‌بندی طیف‌هایی که دارای رفتار یکسانی‌اند «طبقه‌بندی اطلاعات ماهواره‌ای» گفته می‌شود. در این فرایند، برچسب خاصی به هریک از

به‌طور کلی، لندست-۸ معمولاً برای تحقیق و پروژه‌های نیازمند دقت بیشتر و داده‌های تاریخی مناسب‌تر است اما سنتینل گزینه بهتری در زمینه کاربردهای نیازمند اسکن‌های تکراری سریع‌تر و همچنین نظارت و مدیریت منابع طبیعی است. از این رو در این پژوهش، به‌منظور ارزیابی دقت عملکرد الگوریتم‌های طبقه‌بندی، از تصویر ماهواره لندست استفاده شده است.

۲-۲-۲- پیش‌پردازش تصاویر ماهواره‌ای

در بخش پیش‌پردازش تصاویر، تصویر دریافت‌شده با ترکیب مجموعه باندهای ۱، ۲، ۳، ۴، ۵ و ۷ به دست آمد (با توجه به حرارتی بودن باند ۶، از آن استفاده نشده است). از آنجاکه تصویر لندست-۸ از سوی سازمان زمین‌شناسی ایالات متحد به صورت زمین‌مرجع‌شده ارائه شده است، به تصحیح هندسی نیازی ندارد. در ادامه، برای تشخیص عوارض، از ترکیب رنگی کاذب برای نمایش تصاویر استفاده شده است. به‌منظور نمایش



شکل ۲. فرایند تحقیق

ماهواره‌ای‌اند. ماشین بردار پشتیبان، به‌منزله الگوریتمی مبتنی‌بر حاشیه، در تلاش است تا بهترین مرز تفکیک‌کننده را بین کلاس‌ها پیدا کند. جنگل تصادفی الگوریتمی مبتنی‌بر مجموعه است که از چندین درخت تصمیم برای بهبود دقت و کاهش احتمال بیش‌برازش استفاده می‌کند و الگوریتم بیشترین احتمال نیز روشی آماری است که براساس توزیع‌های احتمالی عمل می‌کند. این تنوع مقایسه عملکرد الگوریتم‌ها را، در شرایط متفاوت، امکان‌پذیر می‌سازد (Talukdar et al., 2020; Marudi et al., 2024; Ganesh et al., 2023).

– کارایی در مسائل گوناگون: این الگوریتم‌ها کاربرد گسترده‌ای در مسائل گوناگون طبقه‌بندی و پیش‌بینی دارند و در بسیاری از مطالعات، عملکرد خوبی نشان داده‌اند. انتخاب این الگوریتم‌ها دستیابی به نتایج قابل‌مقایسه و معتبرتری را فراهم می‌آورد. – سازگاری با داده‌های پیچیده: جنگل تصادفی و ماشین بردار پشتیبان می‌توانند از داده‌های پیچیده و غیرخطی استفاده‌ای بهینه کنند. این ویژگی‌ها، به‌خصوص در زمینه‌های جغرافیایی و محیطی که داده‌ها ممکن است دارای الگوهای پیچیده باشند، اهمیت دارد.

– تفسیرپذیری: الگوریتم‌های مورد مطالعه تفسیرپذیری چشمگیری دارند و می‌توانند به کاربران، در درک شیوه‌های تصمیم‌گیری، کمک کنند. این ویژگی می‌تواند در تحلیل نتایج و ارائه آن‌ها به ذی‌نفعان مفید باشد.

– تجربه‌های پیشین: این الگوریتم‌ها، در بسیاری از پژوهش‌های پیشین، به‌صورت الگوریتم‌های مؤثر در زمینه‌های مشابه شناخته شده‌اند. این نکته می‌تواند سبب افزایش اعتبار نتایج و میزان اعتماد به آن‌ها شود.

پیکسل‌های تصاویر داده می‌شود و پیکسل‌های مشابه، در یک طبقه و با برچسبی مشترک، قرار می‌گیرند. به‌طور کلی، الگوریتم‌های طبقه‌بندی تصاویر ماهواره‌ای در چهار دسته^۱ (نظارت‌شده^۲؛ ۲) نیمه‌نظارت‌شده^۳؛ ۳) بدون نظارت^۴؛ ۴) تقویتی قرار می‌گیرند (Dahlke et al., 2020; Glielmo et al., 2021). یادگیری تحت نظارت، فرایند آموزش مدل‌های یادگیری ماشین با استفاده از داده‌های برچسب‌گذاری‌شده است تا بتوانند پیش‌بینی‌ها یا دسته‌بندی‌های دقیقی انجام دهند (Dahlke et al., 2020). الگوریتم‌های ماشین بردار پشتیبان و جنگل تصادفی که در این پژوهش استفاده شده‌اند، در گروه یادگیری نظارت‌شده قرار دارند. عملکرد ماشین بردار پشتیبان یافتن یک ابرصفحه^۵ است که کلاس‌ها را به بهترین شکل از هم جدا می‌کند (Talukdar et al., 2020; Santarsiero et al., 2022; Chandra & Bedi, 2021). اما الگوریتم جنگل تصادفی شیوه‌ای برای مدل‌سازی تصمیم‌گیری است که با تقسیم داده‌ها به زیرمجموعه‌ها براساس ویژگی‌ها، درختی شامل پرسش‌های باینری ایجاد می‌کند تا به پیش‌بینی یا دسته‌بندی برسد (Moradi & Kasani, 2024; Marudi et al., 2024; Naswin & Wibowo, 2023; Jiao et al., 2020). الگوریتم حداکثر احتمال، به‌منزله الگوریتمی تحت نظارت، روشی برای برآورد پارامترهای مدل آماری است و با انتخاب مقادیر پارامترها، احتمال مشاهده داده‌های موجود را به حداکثر می‌رساند (Ganesh et al., 2023; Yimer et al., 2024). الگوریتم‌های یادگیری بدون نظارت، فرایند آموزش مدل‌های یادگیری ماشین با استفاده از داده‌های بدون برچسب است تا الگوها و ساختارهای پنهان در داده‌ها شناسایی شود (Sakai et al., 2017; Dahlke et al., 2020; Gan et al., 2013).

الگوریتم‌های ماشین بردار پشتیبان، جنگل تصادفی و بیشترین احتمال به‌این‌دلایل، در پژوهش حاضر انتخاب شده‌اند:

– تنوع در رویکردها: این سه الگوریتم نمایانگر رویکردهای متفاوت در طبقه‌بندی تصاویر

1. Supervised Learning
2. Semi-Supervised Learning
3. Unsupervised Learning
4. Reinforcement Learning
5. Hyperplane

سادگی زیادی را درمورد استقلال ویژگی‌ها ایجاد می‌کند که درموارد واقعی، ممکن است نادرست باشد؛ اما این ویژگی معمولاً در داده‌های جغرافیایی برآورد نمی‌شود (Ganesh et al., 2023; Yimer et al., 2024). بنابراین الگوریتم‌های منتخب گزینه‌های مناسبی برای طبقه‌بندی منطقه کلان‌شهری مورد اشاره 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 تمامی طبقات انتخاب شده است (۷۵٪؛ ۲) گروه آزمایش که حدوداً ۵۰۰ نقطه، به‌منزله نمونه آموزشی، وارد ماشین می‌شوند، برچسب می‌خورند و آموزش می‌بینند (۲۵٪). با توجه به دقت تصاویر ماهواره‌های لندست-۸ (۳۰ متر) که از سال ۲۰۲۰ تهیه شده است و ویژگی‌های منطقه کلان‌شهری مشهد، نمونه‌های منتخب در پنج طبقه (۱) مناطق ساخته‌شده؛ (۲) اراضی بایر؛ (۳) مناطق کوهستانی؛ (۴) فضاهای سبز؛ (۵) پهنه‌های آبی قرار گرفته‌اند.

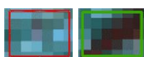
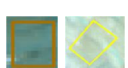
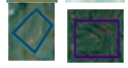
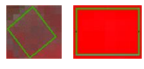
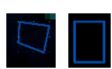
۴-۲-۲- تنظیم هایپرپارامترها

در این بخش از پژوهش، تنظیم هایپرپارامترهای مدل مطرح شده است. هایپرپارامترها مقادیر قابل تنظیمی‌اند که قبل از آموزش مدل‌های یادگیری ماشین، تعیین می‌شوند و در عملکرد و دقت مدل تأثیر می‌گذارند. برخلاف پارامترهای داخلی مدل که در خلال فرایند آموزش به دست می‌آیند، هایپرپارامترها باید به‌صورت دستی یا با استفاده از تکنیک‌های بهینه‌سازی تنظیم شوند. تنظیم مناسب هایپرپارامترها می‌تواند به بهبود تعمیم‌پذیری و کارایی مدل کمک کند (Alonso-Sarría et al., 2024). در این پژوهش پس از انتخاب نمونه‌ها، هایپرپارامترها با تکنیک بهینه‌سازی تنظیم شده است؛ بدین معنی که نرم‌افزار، با توجه به تعداد نمونه‌ها و پیچیدگی آن‌ها، مقادیر هایپرپارامترها را به‌صورت خودکار تنظیم می‌کند. این مقادیر را پس از صحت‌سنجی مدل، می‌توان ارزیابی کرد و به‌صورت دستی تغییر داد.

در هر الگوریتم طبقه‌بندی، معایب خاصی نیز به چشم می‌خورد؛ بنابراین باید الگوریتم‌ها با دقت انتخاب شوند و در این فرایند، ویژگی‌های تحقیق و محدوده مورد مطالعه نیز مد نظر قرار گیرد. الگوریتم ماشین بردار پشتیبان ممکن است، در مواجهه با داده‌های بزرگ و دارای ابعاد گسترده، با مشکلاتی در زمان آموزش و پیش‌بینی روبه‌رو شود و پیچیدگی محاسباتی بسیاری در پی داشته باشد (Talukdar et al., 2020; Santarsiero et al., 2022; Chandra & Bedi, 2021). همچنین امکان دارد جنگل تصادفی، در کار با داده‌های نامتعادل، کارایی مطلوبی نداشته باشد و به مصرف بالای حافظه نیاز داشته باشد (Piryonesi, 2019; Louppe, 2014). باین حال، با توجه به محدودیت‌هایی در وسعت منطقه مورد مطالعه در این پژوهش، معایب یادشده در اینجا بروز نخواهد کرد. علاوه‌براین، الگوریتم بیشترین احتمال سبب می‌شود استقلال ویژگی‌ها بسیار ساده فرض شود که ممکن است، در شرایط واقعی، نادرست باشد؛ هرچند این ویژگی معمولاً در داده‌های جغرافیایی برآورد نمی‌شود (Ganesh et al., 2023; Yimer et al., 2024). به‌همین دلیل الگوریتم‌های انتخابی گزینه‌های مناسبی برای طبقه‌بندی منطقه کلان‌شهری مشهد به شمار می‌آیند.

هریک از الگوریتم‌های طبقه‌بندی معایبی نیز دارد؛ بنابراین انتخاب الگوریتم‌ها باید با احتیاط همراه باشد و در این فرایند، ویژگی‌های تحقیق و محدوده مورد مطالعه مد نظر قرار گیرد. الگوریتم ماشین بردار پشتیبان ممکن است، با داده‌های بزرگ و دارای ابعاد بالا، به‌راحتی دچار مشکلات زمان آموزش و پیش‌بینی و در پی آن، پیچیدگی محاسباتی شود (Talukdar et al., 2020; Santarsiero et al., 2022; Chandra & Bedi, 2021). جنگل تصادفی نیز احتمال دارد، با داده‌های نامتعادل، به‌درستی عمل نکند و افزون‌بر آن به مصرف حافظه معتنابهی نیاز دارد (Piryonesi, 2019; Louppe, 2014). اما از آنجاکه وسعت محدوده مورد مطالعه در این پژوهش زیاد نیست، ایرادهای بیان‌شده در این پژوهش بروز نمی‌یابند. همچنین الگوریتم بیشترین احتمال فرض‌های

جدول ۱. کلاس‌های طبقه‌بندی

نام طبقه	توضیح	مثالی از نمونه انتخابی طبقه
مناطق ساخته شده	مناطق مسکونی، تجاری، صنعتی، شبکه راه‌ها و هر نوع تأسیسات و تجهیزات شهری	
اراضی بایر	اراضی دارای 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	سدها، باتلاق‌ها، رودخانه‌ها، کانال‌ها، استخرهای روباز	

در الگوریتم ماشین بردار پشتیبان، از تابع کرنل گوسی^۱ به همراه هایپرپارامترهای مقدار گامای ۳۲ و C برابر با ۸۱۹۲ استفاده شده است. مقدار بزرگ گاما (۳۲) به این معنی است که تابع کرنل با دقت بیشتری به داده‌های آموزشی نزدیک خواهد شد و درمورد پیچیدگی‌های موجود در توزیع داده‌ها، کارایی بیشتری دارد. این نکته برای طبقه‌بندی دقیق کاربری اراضی که ممکن است الگوهای بسیار متفاوتی داشته باشند، ضروری است. ازسوی دیگر، مقدار C برابر با ۸۱۹۲ نشان‌دهنده تمایل به حداقل نگه داشتن خطا در داده‌های آموزشی بوده و گویای مرز تفکیکی قوی و دقیقی است. این نکته، به موازات توانمندی‌های مدل، در تفکیک دقیق کلاس‌ها کمک می‌کند. به‌طور کلی، این ترکیب از هایپرپارامترها در راستای به حداکثر رساندن کارایی و دقت طبقه‌بندی مدل در منطقه‌ای پیچیده، مانند مشهد، تنظیم شده‌اند.

در الگوریتم جنگل تصادفی، حداکثر تعداد درختان^۲ برابر با ۵۰ و حداکثر عمق آن‌ها^۳ ۳۰ تنظیم شده و حداکثر تعداد نمونه‌ها برای هر کلاس نیز برابر با ۱۰۰۰ تعیین شده است. این مقادیر از هایپرپارامترها به دلیل ایجاد توازن مناسب میان دقت و کارایی مدل در تحلیل داده‌های جغرافیایی در نظر گرفته شده است.

تعداد ۵۰ درخت تنوع کافی را در اختیار قرار می‌دهد تا ویژگی‌های متفاوت داده به‌خوبی شناسایی و از خطر بیش‌برازش جلوگیری شود. عمق حداکثر ۳۰، برای درختان، توانایی مدل را در یادگیری نازک‌کاری‌های پیچیده‌تری که ممکن است در داده‌های جغرافیایی وجود داشته باشد، افزایش می‌دهد. همچنین محدود کردن حداکثر تعداد نمونه‌ها به ۱۰۰۰، در هر کلاس، به بهینه‌سازی اندازه هر زیرمجموعه آموزشی یاری می‌رساند 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 و برچسب زدن آن‌ها در قالب طبقات تعریف شده و تنظیم هایپرپارامترها، مرحله اصلی مدل، یعنی پیش‌بینی سایر پیکسل‌های برچسب‌نخورده 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 مورد نظر و نمونه‌های برداری نقطه‌ای به ماشین، تصویر طبقه‌بندی‌شده کاربری اراضی به‌منزله خروجی مدل در نظر گرفته می‌شود.

1. RBF

2. MaxNumTrees

3. MaxTreeDepth

۲-۲-۵- ارزیابی صحت طبقه‌بندی مدل

برای ارزیابی دقت مدل، باید عملکرد مدل ایجادشده در برابر مجموعه داده‌هایی را که قبلاً وارد مدل نشده است، بررسی کنیم. ماشین، درواقع، الگوها و ویژگی‌های گوناگونی را از داده‌هایی که به آن آموزش داده شده است، فرامی‌گیرد و خود را برای تصمیم‌گیری‌هایی در زمینه‌های گوناگون، مانند شناسایی، طبقه‌بندی یا پیش‌بینی داده‌های جدید آموزش می‌دهد. برای بررسی دقیق اینکه ماشین چگونه قادر به اتخاذ این تصمیم‌هاست، پیش‌بینی‌ها را روی داده‌های آموزش داده‌شده می‌آزمایند. برای این کار، ابتدا روی داده‌های آموزش داده‌شده کار می‌کنیم و پس از آموزش کافی مدل، از آن برای آزمایش براساس داده‌ها استفاده می‌کنیم.

برای ارزیابی صحت طبقه‌بندی، دو روش وجود دارد:

۱) برآورد ماتریس اختلاط^۱ (خطا)؛

۲) محاسبه ضریب توافق کاپا^۲.

۱) برآورد ماتریس اختلاط

این روش یکی از متداول‌ترین شیوه‌های ارزیابی صحت طبقه‌بندی است که با استفاده از آن می‌توان رابطه بین داده‌های مرجع (حقایق زمینی) و نتایج طبقه‌بندی خودکار را مقایسه کرد. در این مرحله، نقاط آزمایشی طبقه‌بندی شده را با تصویر ماهواره‌ای و تصاویر گوگل ارث، و همچنین طبقه‌مدل‌سازی شده را با کاربری واقعی مقایسه می‌کنیم و سپس صحت داشتن یا صحت نداشتن طبقه نقطه آزمایشی را مشخص می‌نماییم. این ماتریس، با ساختاری جدولی، نشان می‌دهد که چه تعداد از رده‌ها به اشتباه به جای رده دیگری قرار گرفته‌اند. در یادگیری بدون نظارت، به آن «ماتریس تطابق»^۳ هم گفته می‌شود.

در هر طبقه‌بندی برای هر نمونه داده، ممکن است یکی از این چهار حالت رخ دهد:

- نمونه عضو دسته مثبت باشد و عضو همین کلاس تشخیص داده شود (مثبت صحیح)^۴؛

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

- نمونه عضو کلاس منفی باشد و عضو همین کلاس تشخیص داده شود (منفی صحیح)^۶؛

- نمونه عضو کلاس منفی باشد و عضو کلاس مثبت تشخیص داده شود (مثبت کاذب)^۷.

پس از اجرای الگوریتم دسته‌بندی، با توجه به توضیحات و تعریف‌های بیان شده، می‌توان عملکرد یک طبقه‌بند را به کمک جدول ۲ بررسی کرد (Arias-Duart et al., 2023).

جدول ۲. حالت‌های گوناگون سلول‌های طبقه‌بندی شده

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

این جدول را اصطلاحاً «ماتریس درهم‌ریختگی» می‌گویند. جدول یا ماتریس درهم‌ریختگی نتایج طبقه‌بندی را براساس اطلاعات واقعی موجود، نمایش می‌دهد. حال، براساس این مقادیر، می‌توان معیارهای گوناگون ارزیابی و اندازه‌گیری دقت را تعریف کرد. معیار صحت^۸ متداول‌ترین، اساسی‌ترین و ساده‌ترین معیار اندازه‌گیری کیفیت دسته‌بندهاست و عبارت است از میزان تشخیص صحیح دسته‌بند در مجموع دو دسته. این معیار درواقع نشان‌دهنده مقدار الگوهای است که درست تشخیص داده شده‌اند و بر مبنای ماتریس ارائه شده در بالا، به صورت رابطه (۱) تعریف می‌شود (Yin et al., 2019).

$$\text{Accuracy} = (TP + TN) / (TP + FN + FP + TN)$$

رابطه (۱). معیار صحت

1. Confusion Matrix

2. Cohen's Kappa Coefficient

3. Matching Matrix

4. True Positive

5. False Negative (FN)

6. True Negative

7. False Positive (FP)

8. Accuracy

ضرب User's accuracy

این ضرب موارد مثبت کاذب را نشان می‌دهد. در این موارد، پیکسل‌ها به اشتباه به منزله کلاسی شناخته شده طبقه‌بندی می‌شوند؛ در حالی که باید به منزله چیز دیگری طبقه‌بندی شوند. این ضرب را «خطای کمیسیون» یا «خطای نوع اول» نیز می‌نامند. داده‌های محاسبه این میزان خطا از ردیف‌های ماتریس اختلاط خوانده می‌شود.

ضرب Producer's accuracy

ضرب موارد منفی کاذب است که در آن پیکسل‌های کلاسی شناخته شده در دسته‌ای غیر از آن کلاس طبقه‌بندی می‌شوند. این ضرب با عنوان «خطاهای حذف» یا «خطای نوع دوم» نیز شناخته می‌شود. داده‌های محاسبه این میزان خطا در ستون‌های جدول خوانده می‌شود.

۲) ضرب توافق کاپا

ضرب کاپا تکنیک چندمتغیره گسسته‌ای است که هرگاه یک ماتریس خطا تفاوت معناداری با ماتریس دیگر داشته باشد، در ارزیابی دقت برای تصمیم‌گیری آماری به کار می‌رود. این ضرب عددی بین ۱- تا ۱+ است؛ هرچه به ۱+ نزدیک‌تر باشد، به توافق بیشتر بین مقیاس‌ها یا

ارزیاب‌ها اشاره دارد و هرچه به ۱- نزدیک‌تر باشد، توافق کمتری را بین آن‌ها نشان می‌دهد. از طرفی، ضرب توافق کاپای برابر صفر نشان‌دهنده وجود نداشتن توافق کامل است (Nti et al., 2023).

جدول ۳ دربردارنده دستورالعمل‌هایی است که لندیس و کوچ^۱ (۱۹۷۷)، فلایس^۲ و همکاران (۱۹۷۱) و بلند و آلتمن^۳ (۱۹۹۵) برای ارزیابی سطح توافق بیان کرده‌اند.

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

رابطه (۲). ضرب کاپا

مقدار PA_o نمایانگر میزان توافق دو ارزیاب است و مقدار PA_E نیز نمایانگر میزان توافق مورد انتظار است.

۳- یافته‌های تحقیق**۳-۱- طبقه‌بندی 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 قرار می‌گیرند؛ به گونه‌ای که در هر سه الگوریتم، اراضی بایر با پوشش بیش از ۴۰٪ از محدوده، کاربری غالب در هر سه الگوی طبقه‌بندی این منطقه به شمار می‌رود و کاربری پهنه‌های آبی نیز که در هر سه الگوریتم کمتر

جدول ۳. آستانه‌های قدرت ضرب توافق کاپا

مقیاس معیار بلند و آلتمن		مقیاس معیار فلایس		مقیاس معیار لندیس و کوچ	
ناسازگار	< ۰/۲	ناسازگار	< ۰/۲	ضعیف	< ۰/۴
ضعیف	۰/۴۰-۰/۲۱	ضعیف	۰/۴۰-۰/۲۱	متوسط تا خوب	۰/۴۰-۰/۷۵
نسبتاً خوب	۰/۶۰-۰/۴۱	نسبتاً خوب	۰/۶۰-۰/۴۱	عالی	بیشتر از ۰/۷۵
خوب	۰/۸۰-۰/۶۱	محکم	۰/۸۰-۰/۶۱		
بسیار خوب	۱۰۰-۰/۸۱	تقریباً کامل	۱۰۰-۰/۸۱		

1. Landis & Koch
2. Fleiss
3. Bland & Altman

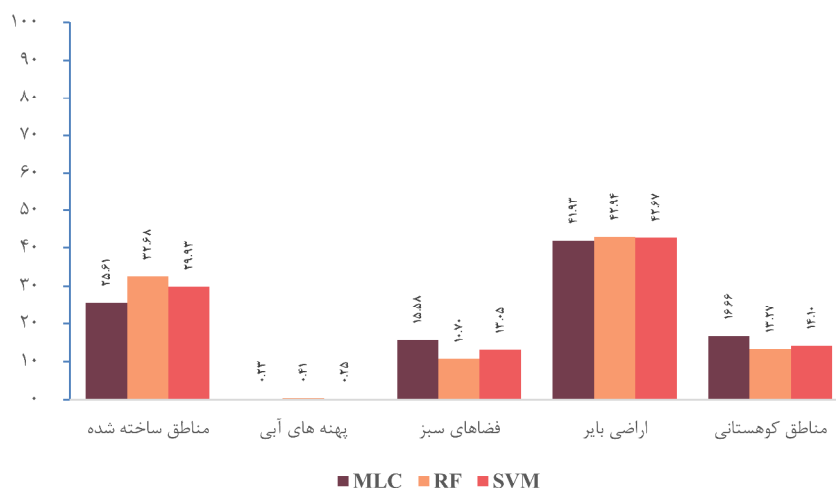
تصادفی، بیشترین میزان ($۰.۳۲/۶۸$) و در الگوریتم بیشترین احتمال، کمترین مقدار ($۰.۲۵/۶۱$) را نشان می‌دهد. درمورد فضاهای سبز نیز، بین سه الگوریتم، تفاوت بسیاری دیده می‌شود اما، برخلاف طبقه مناطق ساخته‌شده، این طبقه در الگوریتم بیشترین احتمال، بیشترین میزان ($۰.۱۵/۵۸$) و در الگوریتم جنگل تصادفی، کمترین مقدار ($۰.۱۰/۷۰$) را داراست.

از ۱% از محدوده را تشکیل می‌دهد، کم‌تعدادترین کاربری در این ناحیه محسوب می‌شود. در جدول ۴، درصد سهم هر کلاس را با توجه به کل پوشش زمین در منطقه مورد مطالعه، برای هر طبقه‌بندی، نشان داده‌ایم. با توجه به این جدول و شکل ۳، مساحت مناطق ساخته‌شده در الگوریتم‌های مورد بررسی تفاوت زیادی را بروز داده‌اند؛ به‌گونه‌ای که این طبقه، در الگوریتم جنگل

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

طبقات کاربری	مناطق ساخته‌شده		پهنه‌های آبی		فضاهای سبز		اراضی بایر		مناطق کوهستانی		مساحت کل (هکتار)
	درصد	هکتار	درصد	هکتار	درصد	هکتار	درصد	هکتار	درصد	هکتار	
بیشترین احتمال (MLC)	۴۶۴۵۶/۸۵	۲۵/۶۱	۴۱۷/۵۱	۰/۲۳	۲۸۲۶۵/۱۲	۱۵/۵۸	۷۶۰۵۶/۷۷	۴۱/۹۳	۳۰۲۱۴/۱۴	۱۶/۶۶	
جنگل تصادفی (RF)	۵۹۲۸۰/۹۳	۳۲/۶۸	۷۳۹/۷۸	۰/۴۱	۱۹۴۱۴/۴۶	۱۰/۷۰	۷۷۸۹۸/۲۱	۴۲/۹۴	۲۴۰۷۶/۶۳	۱۳/۲۷	۱۸۱۴۱۰/۴۰
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 گوناگون کاربری اراضی در منطقه کلان‌شهری مشهد، طی سال ۲۰۲۰

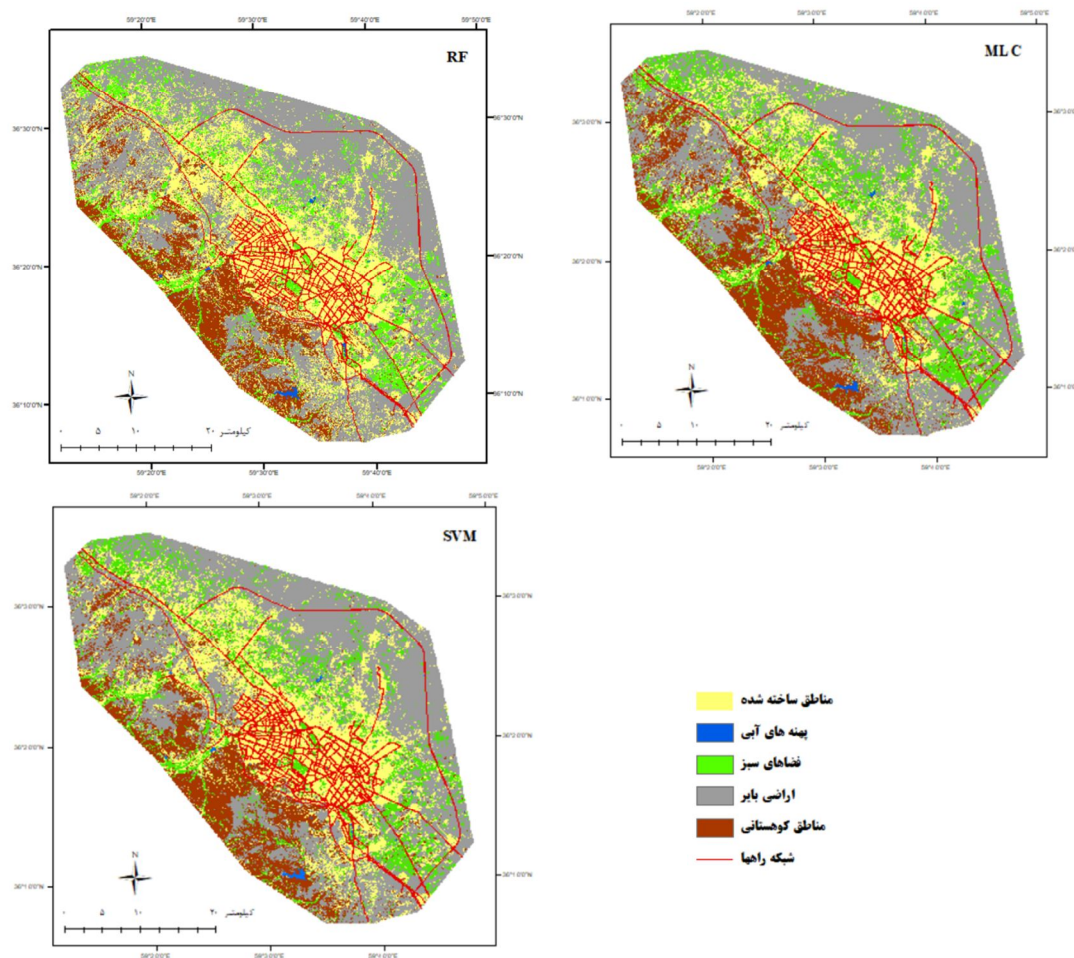
دقت بیشتری طبقه‌بندی می‌شوند زیرا همه طبقه‌بندی‌کننده‌ها دارای مناطق کاملاً یکنواخت، با ضریب تغییرات بسیار اندک‌اند. در مقابل، پهنه‌های آبی و فضاهای سبز به‌خوبی طبقه‌بندی نشده‌اند زیرا تمامی الگوریتم‌ها سهم وسعت طبقات مذکور را با تفاوت‌های شایان توجهی در نظر گرفته‌اند.

۲-۳- اعتبارسنجی طبقه‌بندی‌ها

دقت کلی هریک از الگوریتم‌ها، با استفاده از ماتریس اختلاط و ضریب کاپا (K)، در جدول‌های ۵ تا ۷ بیان شده است. طبقه‌بندی‌کننده SVM به‌منزله مدل بسیار

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 طبقه‌بندی‌شده نشان می‌دهد که طبقه مناطق ساخته‌شده، در الگوریتم جنگل تصادفی، بیش از سایر الگوریتم‌ها در میان مناطق کوهستانی و فضاهای سبز پیش رفته است.

انحراف معیار (SD) و ضریب تغییرات درصد سهم مساحت در کلاس‌های LULC از طریق الگوریتم‌های گوناگون، نیز در جدول ۴ ارائه شده است. ضرایب محاسبه‌شده به‌وضوح بیان می‌کند که اراضی بایر با



شکل ۴. طبقه‌بندی کاربری اراضی منطقه کلان‌شهری مشهد با الگوریتم‌های جنگل تصادفی (RF)، ماشین بردار پشتیبان (SVM) و بیشترین شباهت (MLC)

دقیق LULC، با ضریب کاپای ۰/۹۳، در بین همه طبقه‌بندی‌کننده‌ها شناسایی شده است؛ از این‌رو الگوریتم یادشده، بین نقشه طبقه‌بندی‌شده LULC و واقعیت زمینی، تطابق بسیار خوبی را نشان می‌دهد.

ضریب کاپا در مدل RF برابر با ۰/۸۸ و در مدل MLC برابر با ۰/۸ است؛ بنابراین همه مدل‌ها را می‌توان مفید دانست اما الگوریتم SVM را می‌توان بهترین طبقه‌بندی‌کننده مناسب LULC توصیه کرد.

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

کلاس کاربری	مناطق ساخته‌شده	پهنه‌های آبی	فضاهای سبز	اراضی بایر	مناطق کوهستانی	مجموع	U-Accuracy	Kappa
مناطق ساخته‌شده	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	۰/۹۶۳۹۶۴	۰/۷۱۷۳۹۱۳	۰	۰/۹۱۷۳۲۲۸۳۵	۰
Kappa	۰	۰	۰	۰	۰	۰	۰	۰/۸۸۱۷۶۶

منبع: یافته‌های تحقیق

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

کلاس کاربری	مناطق ساخته‌شده	پهنه‌های آبی	فضاهای سبز	اراضی بایر	مناطق کوهستانی	مجموع	U-Accuracy	Kappa
مناطق ساخته‌شده	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	۰/۹۴۱۹۶۴۳	۱	۰	۰/۹۵۶۷۷۸	۰
Kappa	۰	۰	۰	۰	۰	۰	۰	۰/۹۳۷۹۵

منبع: یافته‌های تحقیق

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

کلاس کاربری	مناطق ساخته‌شده	پهنه‌های آبی	فضاهای سبز	اراضی بایر	مناطق کوهستانی	مجموع	U-Accuracy	Kappa
مناطق ساخته‌شده	۱۱۱	۰	۷	۸	۲	۱۲۸	۰/۸۶۷۱۸۷۵	۰
پهنه‌های آبی	۱	۹	۰	۰	۰	۱۰	۰/۹	۰
فضاهای سبز	۰	۰	۷۹	۰	۰	۷۹	۱	۰
اراضی بایر	۲۷	۰	۴	۱۷۸	۰	۲۰۹	۰/۸۵۱۶۷۴۶	۰
مناطق کوهستانی	۷	۰	۰	۱۴	۶۲	۸۳	۰/۷۴۶۹۸۸	۰
مجموع	۱۴۶	۹	۹۰	۲۰۰	۶۴	۵۰۹	۰	۰
P-Accuracy	۰/۷۶۰۲۷۴	۱	۰/۸۷۷۷۷۸	۰/۸۹	۰/۹۶۸۷۵	۰	۰/۸۶۲۴۷۵۴	۰
Kappa	۰	۰	۰	۰	۰	۰	۰	۰/۸۰۸۵۲۴۲

منبع: یافته‌های تحقیق

بررسی ضریب U-Accuracy در الگوریتم جنگل تصادفی نشان می‌دهد که برخی سلول‌ها به اشتباه به جای مناطق ساخته شده برداشت شده‌اند؛ یعنی در طبقه‌بندی این سلول‌ها خطای نوع اول رخ داده و سلول‌هایی که در دسته مناطق ساخته شده قرار نمی‌گیرند، به اشتباه، در این طبقه قرار گرفتند (میزان صحت = 0.762195122). خطای نوع اول در مورد 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)، بیشتر از سایر طبقات است؛ به گونه‌ای که سلول‌های این دو طبقه به درستی تشخیص داده نشده و در دسته کاربری دیگری قرار گرفته‌اند.

در مجموع، کاربری بایر با کمترین خطای نوع اول و دوم، در قیاس با دیگر 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 مورد قبول فراتر نمی‌رود.

بررسی ضریب U-Accuracy در الگوریتم ماشین بردار پشتیبان نشان می‌دهد تعداد سلول‌هایی که به اشتباه به منزله مناطق کوهستانی برداشت شده‌اند بیش از سایر کاربری‌هاست؛ یعنی، در طبقه‌بندی این سلول‌ها، خطای نوع اول رخ داده است و سلول‌هایی که در دسته مناطق کوهستانی قرار نمی‌گیرند، به اشتباه، در این طبقه قرار گرفتند (میزان صحت = 0.8450704). از سویی، بررسی ضریب P-Accuracy نشان می‌دهد که در طبقات متفاوت، خطای نوع دوم به مقدار بسیار ناچیزی دیده می‌شود؛ به گونه‌ای که حداقل این ضریب در اراضی بایر و برابر با 0.9419643 مشاهده می‌شود. در مجموع، خطای نوع اول و دوم در این الگوریتم بسیار ناچیز است.

بررسی ضریب U-Accuracy در الگوریتم بیشترین احتمال نشان می‌دهد که برخی سلول‌ها، به اشتباه، به منزله مناطق کوهستانی شناسایی شده‌اند؛ یعنی در طبقه‌بندی این سلول‌ها خطای نوع اول رخ داده است و سلول‌هایی که نباید در طبقه مناطق کوهستانی واقع شوند، به اشتباه، در این دسته قرار گرفتند (میزان صحت = 0.746988). خطای نوع اول در مورد 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)، بیشتر از سایر طبقات است؛ به گونه‌ای که سلول‌های این طبقه به درستی تشخیص داده نشده و در دسته کاربری دیگری قرار گرفته است. باین حال، با بررسی خطاهای نوع اول و دوم، می‌توان گفت که درصد صحت این خطاها از بازه مورد قبول بیشتر نمی‌شود.

۴- بحث

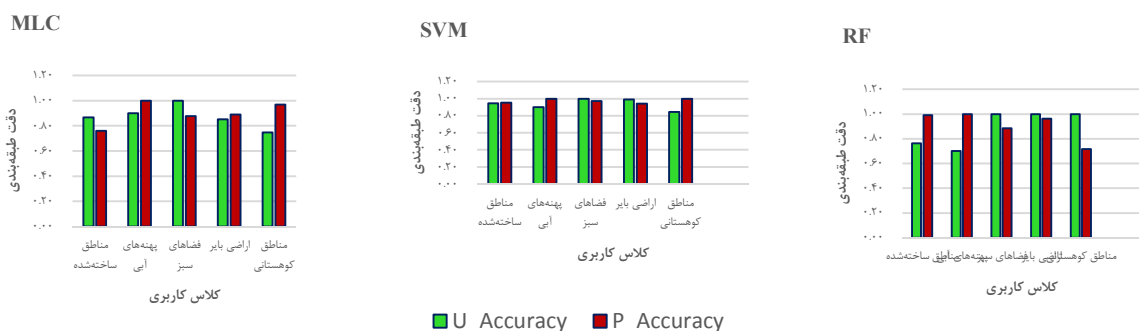
پژوهش‌ها و مطالعات گوناگون نشان داده است که الگوریتم‌های متفاوت در شرایط متفاوت دقت متفاوتی دارند. آلونسو ساریا^۱ و همکاران (۲۰۲۴) به این نتیجه رسیده‌اند که جنگل تصادفی و ماشین بردار پشتیبان تا حدودی عملکرد یکسانی دارند و در مقایسه با الگوریتم بیشترین شباهت، دارای دقت بیشتری هستند. مطالعات گوناگون، از جمله باکاری^۲ و همکاران (۲۰۲۴)، جاسیم^۳ و همکاران (۲۰۲۴)، ساویتا و رشم^۴ (۲۰۲۲) و قدسی و همکاران (۲۰۲۱)، هم‌راستا با نتایج پژوهش پیش رو، اثربخشی SVM را در دستیابی به دقت‌های بالا در طبقه‌بندی کاربری و پوشش زمین نشان داده‌اند.

1. Alonso-Sarria
2. Baccari
3. Jasim
4. Savitha & Reshma

طبقات مناطق ساخته‌شده و کوهستانی، بیش از دو الگوریتم دیگر است؛ به‌گونه‌ای که برخی سلول‌هایی که در طبقه مناطق ساخته‌شده قرار نمی‌گرفتند، به‌منزله مناطق ساخته‌شده شناسایی شدند. این خطا درمورد پهنه‌های آبی نیز صدق می‌کند؛ یعنی سلول‌هایی شناسایی شده‌اند که به‌اشتباه در طبقه پهنه‌های آبی قرار گرفته‌اند. ازطرفی، دقت این الگوریتم در شناسایی مناطق کوهستانی نیز کمتر از سایر الگوریتم‌هاست زیرا نتوانسته است مناطق کوهستانی را به‌دقت شناسایی کند. بررسی فضایی سلول‌های دارای خطا نشان می‌دهد که برخی از سلول‌های متعلق به طبقه مناطق ساخته‌شده به‌رنگ سوخته‌اند و در طیف رنگی مناطق کوهستانی قرار دارند. خطاهای رخ داده به این نوع سلول‌ها اختصاص دارد.

با توجه 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 رخ داده است؛ به‌گونه‌ای که تعدادی از سلول‌هایی که در مناطق کوهستانی نبودند به‌منزله مناطق کوهستانی شناسایی شده‌اند. اما، در مقابل، خطای نوع دوم درمورد این کاربری در الگوریتم SVM کمتر دیده شده است؛ یعنی مناطق کوهستانی واقعی با دقت بسیار بالایی شناسایی شده‌اند.

و به این نتیجه رسیده‌اند که عملکرد الگوریتم ماشین بردار پشتیبان در طبقه‌بندی کاربری اراضی، به‌طور کلی، درقیاس با سایر الگوریتم‌های محبوب مانند جنگل تصادفی و حداکثر احتمال، برتر است. در مقابل، رازافینیمارو^۱ و همکاران (۲۰۲۲) و ستیانو و سونیوتو^۲ (۲۰۲۴)، برخلاف پژوهش حاضر، بر دقت بیشتر جنگل تصادفی درمقایسه با ماشین بردار پشتیبان تأکید کرده‌اند. به‌طور خلاصه، درحالی که RF و SVM همواره دقت بالایی را در طبقه‌بندی LULC نشان می‌دهد، انتخاب الگوریتم ممکن است با توجه به شرایط خاص و ویژگی‌های داده و محدوده مورد مطالعه انجام شود؛ بنابراین چالش‌های شناسایی کاربری‌ها، در منطقه کلان‌شهری مشهد، ازطریق الگوریتم‌های مورد مطالعه این پژوهش نیز بررسی شده است. برای این کار، کاربری‌های طبقه‌بندی شده با تصاویر گوگل ارث مورد مقایسه قرار گرفته است. در این خصوص، آگاهی و شناخت پژوهشگران درباره محدوده مورد مطالعه نیز نقش بسزایی داشته است. این مقایسه نشان داد شباهت رنگ برخی عوارض و مناطق به خطای طبقه‌بندی منجر می‌شود؛ برای نمونه، برخی از بخش‌های ساخته‌شده به‌منزله مناطق کوهستانی و برخی سلول‌های کوهستانی به‌منزله مناطق ساخته‌شده شناسایی شدند. هر دو نوع اول و دوم خطا در این الگوریتم، درمورد



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

1. Razafinimaro
2. Setyanto & Sunyoto

طبقه‌بندی‌کننده LULC را نشان می‌دهد. تفاوت بین دقت طبقه‌بندی‌کننده‌های مورد استفاده جزئی است اما این تغییرات جزئی اهمیت چشمگیری در زمینه برنامه‌ریزی LULC دارد. از آنجاکه این اختلافات جزئی در کاربری‌های حساسی مانند مناطق ساخته‌شده و فضاهای سبز دیده می‌شود، انتخاب الگوریتمی دارای بیشترین دقت و کمترین خطا بسیار اهمیت می‌یابد. نتایج بررسی ضریب کاپا و تحلیل‌های مبتنی بر ماتریس اختلاط نشان می‌دهد که ماشین بردار پشتیبان بیشترین دقت و کمترین خطا را بین طبقه‌بندی‌کننده‌های مورد مطالعه دارد.

طبق یافته‌های مطالعات متعدد، دقت نقشه‌برداری LULC با زمان و مکان در ارتباط است؛ بنابراین، درمورد تحقیقات آینده، تحلیل دقت طبقه‌بندی‌کننده‌ها برای شرایط مورفولوژیک و ژئومورفیک متفاوت پیشنهاد می‌شود.

۶- منابع

- Abburu, S. & Golla, S.B. 2015, **Satellite Image Classification Methods and Techniques: A Review**, International Journal of Computer Applications, 119(8).
- Aburas, M.M., Ho, Y.M., Ramli, M.F. & Ash'aari, Z.H., 2016, **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, PP. 380-389, <https://doi.org/10.1016/j.jag.2016.07.007>.
- Aljanabi, F., Dedeoğlu, M. & Şeker, C., 2024, **Environmental Monitoring of Land Use/Land Cover by Integrating Remote Sensing and Machine Learning Algorithms**, Journal of Engineering and Sustainable Development, 28(4), PP. 455-466, <https://doi.org/10.31272/jeasd.28.4.4>.
- Francisco Alonso-Sarria, F., Valdivieso-Ros, C., Francisco Gomariz-Castillo, F., 2024, **Analysis of the hyperparameter optimisation of four machine learning satellite imagery classification methods**, Computational Geosciences (2024) 28:551–571 <https://doi.org/10.1007/s10596-024-10285-y>.

در الگوریتم MLC به‌منزله کم‌دقت‌ترین الگوریتم، علاوه‌بر مناطق کوهستانی و مناطق ساخته‌شده، شناسایی فضاهای سبز و اراضی بایر نیز با چالش مواجه بود؛ به‌گونه‌ای که برخی از سلول‌های غیربایر به‌اشتباه در طبقات بایر قرار گرفتند. همچنین این الگوریتم، در شناسایی فضاهای سبز نیز، خطایی بیشتر از سایر الگوریتم‌ها داشته است.

به‌طور کلی، می‌توان گفت کاربری بایر با دقت بیشتر و خطای کمتر و کاربری ساخته‌شده و مناطق کوهستانی با خطای بیشتری شناسایی شده‌اند.

۵- نتیجه‌گیری

این مطالعه به‌منظور بررسی دقت عملکرد پرکاربردترین الگوریتم‌های طبقه‌بندی کاربری اراضی، براساس مشاهدات ماهواره‌ای، انجام شد. هدف این پژوهش پیشنهاد بهترین الگوریتم طبقه‌بندی‌کننده بود.

با بررسی ویژگی‌های گوناگون الگوریتم‌های پرکاربرد در طبقه‌بندی کاربری اراضی، سه الگوریتم نظارت‌شده ماشین بردار پشتیبان، جنگل تصادفی و بیشترین احتمال روی داده‌های لندست-۸ (OLI) برای طبقه‌بندی کاربری اراضی منطقه کلان‌شهری مشهد در سال ۲۰۲۰ انتخاب شدند. صحت طبقه‌بندی کاربری‌های متفاوت با استفاده از ضرایب P-Accuracy و U-Accuracy در ماتریس اختلاط و بررسی ضریب تغییرات سهم هر کاربری در الگوریتم‌های مورد مطالعه ارزیابی شد. نتایج بررسی ضرایب یادشده نشان می‌دهد که به‌طور کلی، صحت طبقه‌بندی دسته‌ها در تمامی الگوریتم‌های مورد مطالعه، در بازه خوب تا عالی قرار می‌گیرد. اما بررسی دقیق‌تر این الگوریتم‌ها مشخص می‌کند بیشترین چالش شناسایی طبقه درمورد مناطق ساخته‌شده، مناطق کوهستانی و فضاهای سبز رخ می‌دهد و شناسایی اراضی بایر با چالش کمتری مواجه است. همچنین بررسی ضریب تغییرات سهم طبقات شناسایی‌شده در الگوریتم‌های گوناگون نشان می‌دهد که فضاهای سبز و پهنه‌های آبی همبستگی کمتری در الگوریتم‌های مورد مطالعه دارند. ضریب کاپا و تحلیل‌های مبتنی بر ماتریس اختلاط نیز تنوع در دقت هر

- Azizi Qalati, S., Rangzen, K., Taghizadeh, A. & Ahmadi, S., 2013, **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(4), PP. 585-596.
- Baccari, N., Hamza, M.H., Slama, T., Sebei, A. & Rebai, N., 2024, **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., 1995, **Multiple Significance Tests: The Bonferroni Method**, BMJ, 310(6973), P. 170, <https://doi.org/10.1136/bmj.310.6973.170>.
- Chandra, M.A. & Bedi, S.S., 2021, **Survey on SVM and their Application in Image Classification**, International Journal of Information Technology, 13(5), PP. 1-11, <https://doi.org/10.1007/s41870-017-0080-1>.
- Dahlke, J., Bogner, K., Mueller, M., Berger, T., Pyka, A. & Ebersberger, B., 2020, **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., 2017, **Four Decades of Urban Sprawl and Population Growth in Teresina, Brazil**, Applied Geography, 79, PP. 73-83, <https://doi.org/10.1016/j.apgeog.2016.12.018>.
- Esmaili, A. & Ashjaei, H., 2019, **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(31), PP. 153-172, DOI: 10.22111/GAIJ.2019.4710.
- Fleiss, J.L., 1971, **Measuring Nominal Scale Agreement among Many Raters**, Psychological Bulletin, 76(5), P. 378, <https://doi.org/10.1037/h0031619>.
- Gan, H., Sang, N., Huang, R., Tong, X. & Dan, Z., 2013, **Using Clustering Analysis to Improve Semi-Supervised Classification**, Neurocomputing, 101, PP. 290-298, <https://doi.org/10.1016/j.neucom.2012.08.020>.
- Ganesh, B., Vincent, S., Pathan, S. & Benitez, S.R.G., 2023, **Utilizing LANDSAT Data and the Maximum Likelihood Classifier for Analysing Land Use Patterns in Shimoga, Karnataka**, Journal of Physics: Conference Series (Vol. 2571, No. 1, P. 012001), IOP Publishing, DOI: 10.1088/1742-6596/2571/1/012001.
- Ghodsi, Z., Kheirkhah Zarkesh, M.M. & Gherzmecheshme, B., 2021, **Comparison of Accuracy of Support Vector Machine and Random Forest Methods in Preparing Land Use Map and Crops, Using Sentinel-2 Multi-Temporal Images**, Iranian Remote Sensing and GIS Journal, 12(4), PP. 73-92, DOI: 10.52547/GISJ.12.4.73.
- Gholamali Fard, M., Jorabian Shoshtri, SH., Kohnouj, H. & Mirzaei, M., 2011, **Modelling Land Use Changes of the Coasts of Mazandaran Province Using LCM in GIS Environment**, Ecology, 38(64), PP. 109-124.
- Giulmo, A., Husic, B.E., Rodriguez, A., Clementi, C., Noé, F. & Laio, A., 2021, **Unsupervised Learning Methods for Molecular Simulation Data**, Chemical Reviews, 121(16), PP. 9722-9758. 10.1021/acs.chemrev.0c01195.
- Hosseini, M., Karimi, M., Saadi Masgari, M. & Heydari, M., 2016, **Design and Implementation of an Integrated System of Urban Land Use Change Modeling**, Geographical Sciences, 16(40), PP. 69-91.
- Jahanbakhshi, F. & Ekhtesasi, M.R., 2019, **Performance Evaluation of Three Image Classification Methods (Random Forest, Support Vector Machine and the Maximum Likelihood) in Land Use Mapping**, Water and Soil Science, Science and Technology of Agriculture and Natural Resources, 22(4), PP. 235-247, SID, <https://sid.ir/paper/394972/en>.
- Jasim, B.S., Jasim, O.Z. & AL-Hameedawi, A.N., 2024, **Evaluating Land Use Land Cover Classification Based on Machine Learning Algorithms**, Engineering and Technology Journal, 42(5), PP. 557-568, <http://doi.org/10.30684/etj.2024.144585.1638>.
- Jiao, S.R., Song, J. & Liu, B., 2020, **A Review of Decision Tree Classification Algorithms for Continuous Variables**, Journal of Physics, Conference Series (Vol. 1651, No. 1, p. 012083), IOP Publishing, DOI: 10.1088/1742-6596/1651/1/012083.

- Landis, J.R. & Koch, G.G., 1977, **An Application of Hierarchical Kappa-Type Statistics in the Assessment of Majority Agreement Among Multiple Observers**, *Biometrics*, 33(2), PP. 363-374, <https://doi.org/10.2307/2529786>.
- Louppe, G., 2014, **Understanding Random Forests**, Cornell University Library, 10.
- Mahmoudzadeh, H., Derakhshani, K. & Momeni, S., 2019, **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 1410**, *Human Geography Research*, 51(4), PP. 90-871, 10.22059/JHGR.2018.230788.1007434.
- Marudi, M., Ben-Gal, I. & Singer, G., 2024, **A Decision Tree-Based Method for Ordinal Classification Problems**, *IJSE Transactions*, 56(9), PP. 960-974, <https://doi.org/10.1080/24725854.2022.2081745>.
- Mashhad Municipality, 2016, **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., 2011, **Analysing Land Cover and Land Use Change Processes at Watershed Level: A Multitemporal Study in the Lake Cuitzeo Watershed, Mexico (1975–2003)**, *Applied Geography*, 31(1), PP. 237-250, <https://doi.org/10.1016/j.apgeog.2010.05.010>.
- Mirakhorlou, M. & Rahimzadegan, S., 2018, **Land use Change Modeling Using Markov-Cellular Automata Model and Multi-Criteria Decision Making in Talar Watershed**, *Scientific Journal of Mapping Sciences and Techniques*, 8(1), PP. 85-99.
- Mohammadi, J. & Mahmoud, A., 2012, **Spatial Analysis and Urban Land Use Planning of Dogonbadan (Gachsaran)**, *Geographical Research*, 27(2), PP. 19-36.
- Moradi, M. & Kasani, E., 2024, **Optimal Classification Using Decision Tree**, *Mathematical Researches*, 10(2), PP. 51-65.
- Naswin, A. & Wibowo, A.P., 2023, **Performance Analysis of the Decision Tree Classification Algorithm on the Pneumonia Dataset**, *International Journal of Artificial Intelligence in Medical Issues*, 1(1), PP. 1-9.
- Nti, I.K., Nyarko-Boateng, O., Adekoya, A.F., Weyori, B.A. & Adjei, H.P., 2023, **Predicting Diabetes Using Cohen's Kappa Blending Ensemble Learning**, *International Journal of Electronic Healthcare*, 13(1), PP. 57-70, <https://doi.org/10.1504/IJEH.2023.128605>.
- Piryonesi, S.M., 2019, **The Application of Data Analytics to Asset Management: Deterioration and Climate Change Adaptation in Ontario Roads**, University of Toronto (Canada).
- Pourkhabbaz, H., Mohammadyari, F., Aqdar, H. & Tavakoli, M., 2015, **An Experimental Approach in Modeling Land Use Changes in Behbahan City Using Multi-Temporal Satellite Images**, *Town & Country Planning*, 7(2), PP. 2008-7047, 10.22059/JTCP.2015.57187.
- Qiang, W. & Zhongli, Z., 2011, **Reinforcement Learning Model, Algorithms and Its Application**, *International Conference on Mechatronic Science, Electric Engineering and Computer (MEC)* (PP. 1143-1146), IEEE 10.1109/MEC.2011.6025669.
- Razafinimaro, A., Hajalalaina, A.R., Rakotonirainy, H. & Zafimarina, R., 2022, **Land Cover Classification Based Optical Satellite Images Using Machine Learning Algorithms**, *International Journal of Advances in Intelligent Informatics*, 8(3), PP. 362-380, <https://doi.org/10.26555/ijain.v8i3.803>.
- Ritsema van Eck, J. & Koomen, E., 2008, **Characterising Urban Concentration and Land-Use Diversity in Simulations of Future Land Use**, *The Annals of Regional Science*, 42, PP. 123-140, <https://doi.org/10.1007/s00168-007-0141-7>.
- Sakai, T., Plessis, M.C., Niu, G. & Sugiyama, M., 2017, **Semi-Supervised Classification Based on Classification from Positive and Unlabeled Data**, *International Conference on Machine Learning* (PP. 2998-3006), PMLR.
- Santarsiero, V., Nolè, G., Lanorte, A., Tucci, B., Cillis, G. & Murgante, B., 2022, **Remote Sensing and Spatial Analysis for Land-Take Assessment in Basilicata Region (Southern Italy)**, *Remote Sensing*, 14(7), P. 1692, <https://doi.org/10.3390/rs14071692>.
- Savitha, C. & Reshma, T., 2022, **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**, International Virtual Conference on Developments and Applications of Geomatics, PP. 379-392, Singapore: Springer Nature Singapore.
- Setyanto, A. & Sunyoto, A., 2024, **Comparative Study SVM and Random Forest Algorithms for the Classification of Terrestrial Visual Rock Types**, IOP Conference Series: Earth and Environmental Science 1357(1), 10.1088/1755-1315/1357/1/012036.
- Shi, S., Zhang, X. & Fan, W., 2019, **Explaining the Predictions of any Image Classifier via Decision Trees**, arXiv preprint arXiv, 1911.01058, <https://doi.org/10.48550/arXiv.1911.01058>.
- Sowmya, D.R., Shenoy, P.D. & Venugopal, K.R., 2017, **Remote Sensing Satellite Image Processing Techniques for Image Classification: A Comprehensive Survey**, International Journal of Computer Applications, 161(11), PP. 24-37.
- Talukdar, S., Singha, P., Mahato, S., Pal, S., Liou, Y.A. & Rahman, A., 2020, **Land-Use Land-Cover Classification by Machine Learning Classifiers for Satellite Observations—A Review**, Remote Sensing, 12(7), P. 1135, <https://doi.org/10.3390/rs12071135>.
- Yimer, S.M., Bouanani, A., Kumar, N., Tischbein, B. & Borgemeister, C., 2024, **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(5), PP. 2180-2199, <https://doi.org/10.1016/j.asr.2024.06.010>.
- Yin, M., Wortman Vaughan, J. & Wallach, H., 2019, **Understanding the Effect of Accuracy on Trust in Machine Learning Models**, Proceedings of the 2019 Chi Conference on Human Factors in Computing Systems (PP. 1-12), <https://doi.org/10.1145/3290605.3300509>.

This Page is Intentionally Left Blank