

Кисіль Тетяна Миколаївна*Державний університет інформаційно-комунікаційних технологій, м. Київ*
ORCID 0000-0002-5123-0768**Халапова Софія Веніамінівна***Державний університет інформаційно-комунікаційних технологій, м. Київ***Жебка Вікторія Вікторівна***Державний університет інформаційно-комунікаційних технологій, м. Київ*
ORCID 0000-0003-4051-1190**Стражніков Андрій Анатолійович***Державний університет інформаційно-комунікаційних технологій, м. Київ*
ORCID 0009-0000-3492-2968

ГІПЕРСПЕКТРАЛЬНИЙ АНАЛІЗ ДЛЯ РОЗПІЗНАВАННЯ ТА КЛАСИФІКАЦІЇ МАТЕРІАЛІВ ЗЕМНОЇ ПОВЕРХНІ

Анотація. В даній статті проаналізовано сучасні підходи до класифікації гіперспектральних зображень, включно з методами машинного навчання, які демонструють високу ефективність. Гіперспектральний аналіз дозволяє точно ідентифікувати матеріали на земній поверхні завдяки спектральним характеристикам об'єктів. Проте класифікація гіперспектральних зображень залишається складною через високу розмірність даних, нестачу позначених зразків і значні обчислювальні витрати.

Авторами розглянуто алгоритми XGBoost та Random Forest, які використовуються для точного розпізнавання матеріалів на основі їхніх спектральних характеристик, що є основою розробки ефективної системи класифікації об'єктів земної поверхні та дозволяє зменшити обчислювальні витрати, зберегти точність розпізнавання. Здійснено порівняльний аналіз продуктивності двох обраних моделей за такими показниками, як Confusion Matrix, Accuracy, Sensitivity, Specificity, Precision та Recall.

Особливу увагу приділено проблемі зменшення розмірності даних, що є необхідним етапом у класифікації гіперспектральних зображень. Використання відповідних методів дозволяє зменшити навантаження на обчислювальні системи та підвищити продуктивність алгоритмів. Окрім того, у дослідженні розглядається адаптивність класифікаційних моделей до різних наборів даних, що є важливим аспектом для їхнього подальшого застосування у реальних умовах.

Результати дослідження показують, що запропонована система класифікації забезпечує точне визначення матеріалів на земній поверхні. Отримані висновки можуть бути застосовані у військовій розвідці для виявлення об'єктів, моніторингу змін навколишнього середовища та розпізнаванні потенційно небезпечних матеріалів. Запропоновані методи та моделі дозволяють підвищити ефективність аналізу гіперспектральних даних, що відкриває нові перспективи для подальших досліджень даної галузі.

Ключові слова: гіперспектральний аналіз, ROSIS, XGBoost, RandomForestClassifier, задача класифікації, розпізнавання зображень.

Kysil Tetiana*State University of Information and Communication Technologies, Kyiv*
ORCID:0000-0002-5123-0768**Halapova Sofia***State University of Information and Communication Technologies, Kyiv*

Zhebka Viktoriia*State University of Information and Communication Technologies, Kyiv*
ORCID 0000-0003-4051-1190**Strazhnikov Andrii***State University of Information and Communication Technologies, Kyiv*
ORCID 0009-0000-3492-2968

HYPERSPECTRAL ANALYSIS FOR RECOGNITION AND CLASSIFICATION OF EARTH'S SURFACE MATERIALS

Abstract. *This paper analyzes modern approaches to hyperspectral image classification, including machine learning methods that demonstrate high efficiency. Hyperspectral analysis enables precise identification of Earth's surface materials by utilizing the spectral characteristics of objects. However, hyperspectral image classification remains challenging due to the high dimensionality of data, the lack of labeled samples, and significant computational costs.*

The authors examine the XGBoost and Random Forest algorithms, which are used for accurate material recognition based on their spectral properties. These methods allow the development of an effective classification system of Earth's surface objects, while reducing computational costs and maintaining recognition accuracy. A comparative performance analysis of the two selected models is conducted based on metrics such as Confusion Matrix, Accuracy, Sensitivity, Specificity, Precision, and Recall.

Particular attention is paid to the problem of dimensionality reduction. Such data compression is a crucial step in hyperspectral image classification. The application of appropriate methods helps reduce the computational complexity and improves algorithm performance. Additionally, the study explores the adaptability of classification models to different datasets, which is essential for their practical application in the real-world scenarios.

The research results show that the proposed classification system ensures accurate identification of Earth's surface materials. The findings can be applied in the military intelligence for object detection, environmental change monitoring, and the recognition of potentially hazardous materials. The proposed methods and models enhance the efficiency of hyperspectral data analysis, providing new opportunities for further research in this field.

Keywords: *Hyperspectral analysis, ROSIS, XGBoost, RandomForestClassifier, classification task, image recognition*

Постановка проблеми. Гіперспектральний аналіз зображень є перспективною технологією, що дозволяє отримувати детальну інформацію про матеріали та об'єкти на земній поверхні шляхом аналізу їхнього спектрального підпису. Гіперспектральні датчики забезпечують високу спектральну роздільну здатність, що дає змогу ідентифікувати матеріали, визначати їхній хімічний склад, текстуру та інші фізичні властивості. Така технологія має широке застосування в екологічному моніторингу, медицині, військовій розвідці та інших галузях.

Однією з головних проблем аналізу гіперспектральних зображень є висока розмірність отримуваних даних та значна подібність спектральних характеристик різних матеріалів, що ускладнює їх класифікацію. Традиційні методи обробки гіперспектральних даних потребують значних обчислювальних ресурсів і не завжди забезпечують достатню точність розпізнавання об'єктів. Це особливо критично у військовій сфері, де точне визначення матеріалів може відігравати вирішальну роль у розвідці, виявленні небезпечних об'єктів та плануванні стратегічних операцій.

Гіперспектральний аналіз відіграє ключову роль у екологічному моніторингу, дозволяючи відстежувати зміни в рослинному покриві, контролювати забруднення водних ресурсів і оцінювати стан ґрунтів. У медицині ця технологія використовується для ранньої діагностики захворювань, зокрема онкологічних, шляхом аналізу змін на клітинному рівні. У військовій та безпековій сфері гіперспектральні камери, встановлені на дронах, використовуються для виявлення мін, прихованих об'єктів та інших загроз.

Дослідження, представлене в цій роботі, спрямоване на розробку ефективного підходу до класифікації матеріалів на основі гіперспектральних зображень, що поєднує методи зменшення розмірності та алгоритми машинного навчання. Це дозволяє оптимізувати обчислювальні ресурси та підвищити точність розпізнавання. Розроблена програма дозволяє автоматизувати процес аналізу спектральних характеристик матеріалів, що може суттєво підвищити ефективність гіперспектрального аналізу в реальних умовах.

Класифікація об'єктів на земній поверхні за допомогою гіперспектрального аналізу залишається актуальною проблемою, оскільки сучасні методи часто стикаються з труднощами, такими як нестача позначених даних, висока розмірність зображень та значні обчислювальні витрати. Актуальність даного дослідження обумовлена необхідністю розробки ефективних методів обробки гіперспектральних даних, що дозволяють значно покращити точність класифікації матеріалів, мінімізувати похибки розпізнавання та розширити можливості практичного застосування гіперспектрального аналізу в різних сферах.

Аналіз останніх досліджень і публікацій. Гіперспектральний аналіз зображень є важливим напрямом сучасних наукових досліджень, що має широке застосування в екології, медицині, військовій сфері та інших галузях. Останні наукові роботи присвячені розробці методів класифікації гіперспектральних зображень [1], спрямованих на підвищення точності розпізнавання матеріалів та оптимізацію обчислювальних ресурсів.

Автори Xiaozhen Wang, Jiahang Liu, Weijian Chi, Weigang Wang та Yue Ni [4] у своїй роботі розглядають сучасні методи класифікації гіперспектральних зображень. Вони підкреслюють, що однією з основних проблем є нестача позначених даних, що значно ускладнює навчання моделей машинного навчання. Незважаючи на розвиток методів аугментації даних та напівконтрольованого навчання, проблема залишається актуальною, що обмежує можливості класифікації в реальних умовах.

У дослідженні Hüseyin Firat та Davut Hanbay [3] запропоновано підхід на основі тривимірної згорткової нейронної мережі (3D CNN) з архітектурою ResNet50. Метод демонструє високу ефективність, однак має низку обмежень, зокрема високу обчислювальну складність і чутливість до параметрів моделі. Це може обмежувати його практичне застосування, особливо у випадках, коли обчислювальні ресурси є обмеженими, наприклад, при роботі з бортовими системами аналізу гіперспектральних зображень.

Порівняльний аналіз чотирьох методів контрольованого навчання для класифікації гіперспектральних зображень проведено у роботі Hasna Nhaila, Asma Elmaizi, Elkebir Sarhrouni та Ahmed Hammouch [7]. Дослідження показало, що методи SVM із RBF ядром та Random Forest демонструють найкращі результати серед розглянутих алгоритмів [10]. Проте автори зазначають, що вибір оптимального методу класифікації значною мірою залежить від особливостей набору даних. Це вказує на необхідність розробки адаптивних алгоритмів, які могли б автоматично підбирати найбільш ефективний метод для конкретного набору гіперспектральних зображень.

Незважаючи на значний прогрес у класифікації гіперспектральних зображень, у цій галузі залишається низка невирішених проблем. Однією з основних є *нестача позначених даних*, оскільки навчання високоточних моделей потребує великих обсягів анотованих зображень, що є трудомістким і дорогим процесом. Крім того, значні труднощі спричиняє висока розмірність гіперспектральних зображень, які містять сотні спектральних каналів. Це ускладнює їхню обробку та потребує застосування ефективних методів зменшення розмірності, які б зберігали важливу інформацію.

Ще однією проблемою є *висока обчислювальна складність* сучасних методів, зокрема використання глибоких нейронних мереж, таких як 3D CNN. Хоча вони забезпечують високу точність класифікації, їхнє навчання потребує значних обчислювальних ресурсів, що може обмежувати практичне застосування. Окрім цього, більшість існуючих моделей мають *низьку універсальність*, оскільки вимагають попереднього налаштування для кожного окремого набору даних, що знижує їхню ефективність у різних умовах.

З огляду на ці проблеми, подальші дослідження зосереджені на розробці адаптивних методів класифікації гіперспектральних зображень, які поєднують алгоритми машинного

навчання з методами зменшення розмірності. Важливим напрямом є також створення універсальних підходів, здатних враховувати специфіку оброблюваних даних і забезпечувати високу точність незалежно від характеристик конкретного набору зображень.

Мета і задачі дослідження. Метою дослідження являється розробка системи класифікації об'єктів на земній поверхні з використанням алгоритмів машинного навчання, що дозволить з високою точністю ідентифікувати матеріали. Для її досягнення необхідно проаналізувати застосування гіперспектральних зображень у класифікаційних задачах, дослідити можливості алгоритмів XGBoost та Random Forest у розпізнаванні матеріалів, а також провести порівняльний аналіз їхньої ефективності за показниками Confusion Matrix, Accuracy, Sensitivity, Specificity, Precision та Recall. Крім того, важливо детально розглянути ключові етапи класифікації матеріалів, описати запропоновану систему та обґрунтувати вибір найбільш ефективної моделі.

Дослідження спрямоване на введення у науковий обіг нових фактів щодо оптимізації методів класифікації гіперспектральних зображень. Зокрема, аналіз ефективності різних моделей дозволить уточнити їхні переваги та обмеження, що може стати основою для подальших досліджень у цій галузі. Запропонована система класифікації сприятиме покращенню точності розпізнавання матеріалів та забезпечить ширші можливості для практичного використання гіперспектрального аналізу в різних сферах.

Результати дослідження. Гіперспектральні датчики забезпечують можливість отримання детальної інформації про зображення, розділяючи спектр на численні вузькі смуги, що дозволяє з високою точністю визначати характеристики об'єктів. Гіперспектральне зображення (HSI) поєднує традиційне зображення з спектроскопічним аналізом, що дає змогу одночасно отримувати як просторову, так і спектральну інформацію. Такий підхід [6] дозволяє ідентифікувати та класифікувати матеріали на основі їх спектрального профілю, що є надзвичайно корисним для різних застосувань, зокрема для дистанційного зондування, екологічного моніторингу, медичної діагностики та військової розвідки.

HSI надає багатошаровий спектральний профіль кожного пікселя, що дозволяє детально аналізувати властивості матеріалів, такі як хімічний склад, текстура та вологість. Завдяки цьому можна відстежувати зміни в рослинності, контролювати забруднення водних ресурсів, а також проводити ранню діагностику патологій, наприклад, раку шкіри. У військових застосуваннях технологія використовується для точного визначення військових цілей, оскільки один піксель може представляти площу близько 1–2 м², що робить її незамінною для виявлення мін і вибухонебезпечних предметів [8].

Дане дослідження ґрунтується на наборі даних, отриманому за допомогою гіперспектрального датчика ROSIS [2] під час польоту над містом Павія на півночі Італії. Сенсор ROSIS, який працює на базі приладу із зарядовим зв'язком (ПЗЗ), охоплює спектр електромагнітного випромінювання від 400 до 2500 нм, що дозволяє проводити точні спектральні вимірювання. Зображення з набору даних мають розмір 610×340 пікселів, кожен з яких представлений 103 спектральними каналами в діапазоні 430–860 нм, а просторове розділення становить 1,3 м на піксель. Датасет [5] містить 9 класів об'єктів, серед яких асфальт, гравій, дерева, бетон, цегла та інші матеріали.

Для ефективного розпізнавання спектрів зображень модель XGBoost потребує оптимізації вхідного набору даних [9]. Аналіз головних компонент (PCA) допомагає зменшити розмірність даних із збереженням важливої інформації, необхідної для розпізнавання спектрів. Основна мета PCA — виділити ключові характеристики даних і виявити їхні взаємозв'язки, що дозволяє отримати компактний, але інформативний набір ознак для навчання моделі. PCA досягає цього шляхом перетворення даних у нову систему координат, де головні компоненти є ортогональними і відображають напрямки максимальної дисперсії. Для визначення цих компонентів метод знаходить власні вектори коваріаційної матриці, що забезпечує ефективне скорочення розмірності без суттєвих втрат у точності.

Після нормалізації та зменшення розмірності даних виконується класифікація за допомогою XGBoost. Цей алгоритм представляє собою оптимізовану розподілену бібліотеку градієнтного бустингу, розроблену для досягнення високої продуктивності, гнучкості та

точності прогнозування. XGBoost використовує ітеративний підхід, де кожне нове дерево рішень будується з урахуванням помилок попередніх прогнозів. На вхід алгоритму подається набір ознак (X) для передбачення та цільові значення (Y), які модель має правильно визначити.

Робота алгоритму починається з початкового прогнозу, після чого він аналізує помилки й використовує їх для покращення наступних дерев. На кожному кроці XGBoost оцінює якість вузлів за допомогою метрики подібності (Similarity Score), що визначає, наскільки ефективно вузол узгоджується з поточними залишковими помилками. Значення цього критерію розраховується як відношення квадрата суми помилок до виразу дисперсії $P(1-P)$ із додаванням регуляризаційного параметра. Чим вищий показник, тим краще вузол сприяє зменшенню похибки, що дозволяє алгоритму оптимально розбивати дерево та поступово підвищувати точність прогнозів. У алгоритмі XGBoost прогноз для нового зразка розраховується шляхом підсумовування внеску всіх дерев рішень у модель. Для розрахунку оцінки подібності (Similarity Score, SS) кожного вузла дерева рішень можна визначити за формулою (1). Наразі, корисним наскільки буде поділ вузла для покращення точності прогнозу:

$$SS = \frac{(\sum residuals)^2}{\sum P_{i-1} \cdot (1-P_{i-1}) + \lambda}, \quad (1)$$

де SS - оцінка подібності вузла, $\sum residuals$ - сума залишкових похибок у вузлі, P_{i-1} - ймовірність прогнозу перед поділом, $1-P_{i-1}$ - дисперсія ймовірностей для бінарної класифікації, λ - регуляризаційний параметр.

Цей показник відіграє ключову роль у побудові дерев у XGBoost, дозволяючи алгоритму ефективно вибирати найкращі точки поділу. Після того, як алгоритм XGBoost визначив похибки передбачень, наступним кроком є пошук найкращого способу поділу даних для мінімізації цих похибок. Для цього алгоритм аналізує кожну вхідну ознаку, сортує всі її значення у зростаючому порядку та визначає можливі місця розбиття дерева. Алгоритм переглядає кожну пару сусідніх значень у відсортованому списку, обчислює середнє значення між ними та використовує його як *поріг розгалуження*. Надалі виконується розбиття вибірки на два підмножини:

- ✓ *ліва група* містить значення, які менші або рівні пороговим;
- ✓ *права група* містить значення, які перевищують порогові значення.

На наступному етапі важливим моментом буде визначення оцінки якості розбиття. Для кожного можливого поділу алгоритм обчислює, наскільки якісно покращується модель прогнозу. Це визначається за формулою приросту (2), яка розраховує вигоду від створення нового розгалуження в дереві рішень:

$$Gain = SS_{left\ child} + SS_{right\ child} - SS_{root}, \quad (2)$$

де $SS_{left\ child}$ — оцінка подібності (Similarity Score) для лівого дочірнього вузла, $SS_{right\ child}$ — оцінка подібності для правого дочірнього вузла, SS_{root} — оцінка подібності батьківського вузла перед розбиттям.

Алгоритм перебирає всі можливі варіанти розбиття та вибирає той, який дає найбільший приріст $Gain$, що максимізує точність передбачень. Щоб уникнути надмірного розгалуження дерева, в XGBoost застосовується *усічення (Pruning) вузлів*. Гіперпараметр γ , який використовується в розрахунку (3), визначає мінімальний поріг приросту. Якщо $Gain$ батьківського вузла менший за γ , то цей вузол обрізається:

$$Gain < \gamma \Rightarrow \text{видалення вузла} \quad (3)$$

Це допомагає усунути малозначущі вузли, зменшуючи ймовірність перенавчання заданої моделі. В подальшому виконується обчислення вихідного значення для вузлів. Після визначення оптимального розбиття, XGBoost обчислює вихідне значення (4) кожного вузла:

$$Output\ Value = \frac{\sum residuals}{\sum P_{i-1} \cdot (1-P_{i-1}) + \lambda}, \quad (4)$$

де $\sum residuals$ — сума залишкових похибок (різниця між передбаченими та фактичними значеннями), P_{i-1} — ймовірність передбачення перед поточним поділом, λ — регуляризаційний параметр, який допомагає уникнути перенавчання.

Щоб адаптувати значення для подальшої класифікації, XGBoost використовує перетворення початкової ймовірності (0.5) у логарифмічні шанси (*log-odds*) за формулою (5):

$$\log(odds) = \log\left(\frac{P}{1-P}\right) \quad (5)$$

Отримане значення коригується, додаючи вихідне значення вузла, помножене на коефіцієнт навчання за формулою (6):

$$\log(odds)_{new} = \log(odds)_{old} + Output\ Value \times \eta, \quad (6)$$

де η — коефіцієнт навчання (*learning rate*), який контролює швидкість оновлення прогнозу.

За кожну епоху розрахунку оновлене значення перетворюється на ймовірнісне за формулою (7):

$$P = \frac{1}{1 + e^{-\log(odds)}} \quad (7)$$

Завдяки такому підходу алгоритм XGBoost поступово покращує передбачення, навчаючи послідовність дерев рішень, кожне з яких враховує помилки попередніх моделей. У кінцевому підсумку модель використовує отримані прогнозовані значення для класифікації пікселів гіперспектрального зображення.

Метод XGBoost дозволив отримати класифіковані візуальні результати. Хоча алгоритм успішно розпізнав будівлі, метали та каміння, класифіковане зображення містить значну кількість шумів, що свідчить про невисоку точність моделі. Спостерігається наявність розсіяних пікселів, які, ймовірно, є наслідком помилкової класифікації. Крім того, межі між класами виглядають розмитими порівняно з референсним зображенням, що може вказувати на недостатню якість розподілу класів (рис. 1).

Метод *RandomForestClassifier* було використано як альтернативу *XGBoost* для класифікації зображень. На початковому етапі дані підлягали попередній обробці, зокрема масштабуванню, що забезпечує коректну роботу моделі. Для цього застосовувався метод *StandardScaler*, який приводить кожну ознаку до нульового середнього значення та одиничного стандартного відхилення. Це дозволяє всім ознакам бути в одному масштабі, що покращує стабільність та швидкість навчання моделі.

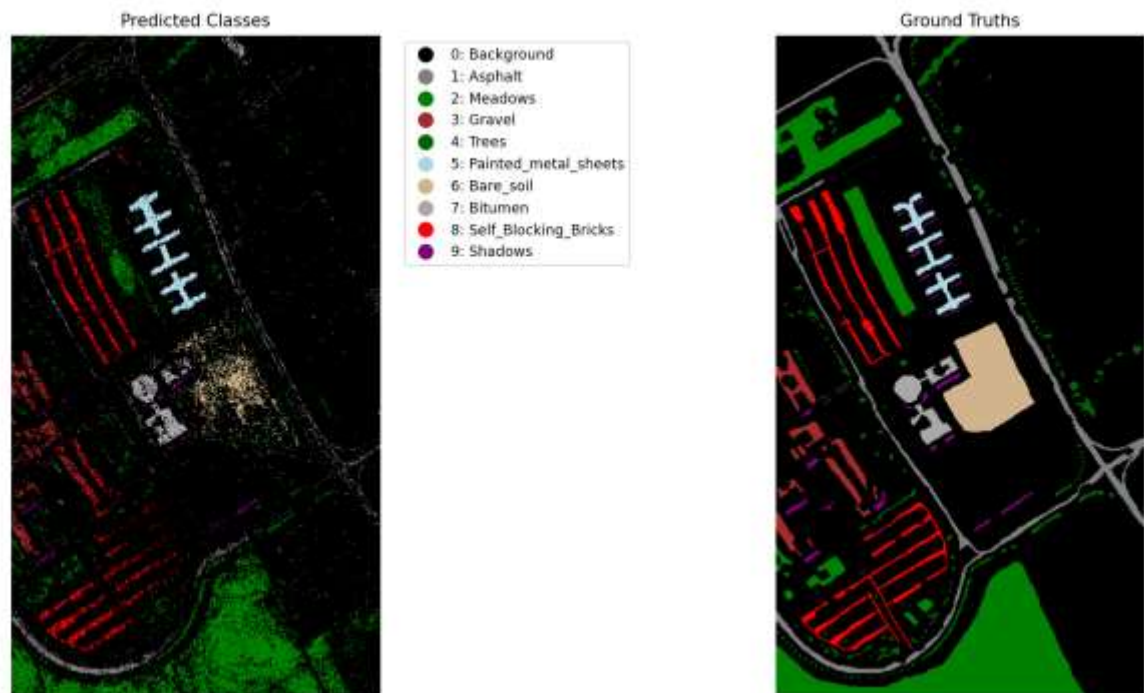


Рис. 1. Результати класифікації вхідних даних за методом XGBoost

Під час масштабування ознак застосовувався процес нормалізації даних, який здійснюється за формулою (8):

$$Z = \frac{X - \mu}{\sigma}, \quad (8)$$

де X – початкове значення ознаки, μ – середнє значення кожного спектрального каналу у тренувальному наборі, σ – стандартне відхилення, Z – стандартизоване значення ознаки.

Після нормалізації дані розбиваються на тренувальну та тестову вибірки. Далі модель *RandomForestClassifier* навчається із застосуванням балансування класів, що забезпечує рівномірний розподіл даних між класами. Алгоритм *Random Forest* базується на ансамблевому методі навчання, який використовує множину дерев рішень (*Decision Trees*) для створення остаточного прогнозу. Його робота починається з формування випадкових підвибірок даних. Для кожного дерева створюється окрема випадкова підмножина ознак та випадковий піднабір навчальних зразків за допомогою методу *Bootstrap Sampling*. Такий підхід дозволяє зменшити кореляцію між деревами та запобігти перенавчанню.

Наступним етапом є навчання окремих дерев рішень. Кожне дерево навчається на власній підвибірці, що робить його унікальним і підвищує стійкість моделі. У процесі навчання використовується певний критерій розбиття, який визначає якість поділу. До таких критеріїв належать *індекс Джині* або *ентропія*, які допомагають обрати оптимальні вузли для подальшого розгалуження дерева. Індекс Джині обчислюється за формулою (9):

$$G = 1 - \sum_{i=1}^n p_i^2, \quad (9)$$

де p_i – ймовірність належності випадкового елемента до i -го класу, n – загальна кількість класів.

Чим менше значення *індексу Джині* (GG), тим чистішим є вузол, що означає ефективніший поділ даних. Після формування та навчання дерев рішень відбувається об'єднання передбачень. Кожне дерево приймає власне рішення щодо класу певного зразка, а остаточне передбачення визначається методом *голосування більшості* (*Majority Voting*), де клас, який отримав найбільшу кількість голосів від дерев, стає остаточним прогнозом та визначається за формулою (10):

$$Y_{final} = \arg \max_k \sum_{i=1}^N I(y_i = k), \quad (10)$$

де Y_{final} – остаточний прогноз, N – загальна кількість дерев лісу, $I(y_i=k)$ – індикаторна функція, яка дорівнює 1, якщо дерево i класифікує об'єкт як клас k , 0 в іншому випадку.



Рис. 2. Результат класифікації вхідних даних за методом *RandomForestClassifier*

Використання випадкових дерев сприяє усуненню шуму в даних і підвищенню точності моделі. Такий підхід допомагає уникнути перенавчання, роблячи модель *стійкішою до шуму* та ефективною навіть при роботі з *високорозмірними та нерозподіленими даними*.

Застосування *RandomForestClassifier* дозволило отримати класифіковане зображення з мінімальною кількістю шумів, що вказує на високу точність моделі. Візуально видно, що

модель забезпечує чіткіші межі між класами порівняно з XGBoost. Завдяки своїй стійкості до шуму метод продемонстрував ефективність у розпізнаванні матеріалів земної поверхні. Однак, як і в будь-якій моделі, були певні недоліки – наприклад, часткове змішування таких класів, як ґрунт та поля (рис. 2).

Метод *Random Forest* показав кращу узагальнюючу здатність завдяки своїй ансамблевій природі. Він дозволив досягти точнішої класифікації матеріалів земної поверхні та показав вищу стійкість до шуму у порівнянні з XGBoost. Навіть без детального аналізу видно, що *RandomForestClassifier* забезпечує більш точну класифікацію матеріалів на земній поверхні.

Порівнюючи методи XGBoost та *RandomForestClassifier* за ключовими метриками, такими як Confusion Matrix, Accuracy, Sensitivity, Specificity, Precision і Recall, було встановлено, що точність XGBoost становить 0.8473, тоді як *RandomForestClassifier* досяг точності 0.9216. Однією з причин цього є вимоги XGBoost до масштабування ознак. Неправильна нормалізація або стандартизація може призвести до втрати інформації, що важлива для класифікації.

Натомість *RandomForestClassifier*, завдяки використанню випадкових дерев рішень, не потребує нормалізації вхідних даних, оскільки працює з кожною ознакою окремо у вузлах дерева. Це дає йому перевагу у випадках, коли ознаки мають різний порядок величини. XGBoost реалізує метод градієнтного бустингу, який поступово покращує передбачення шляхом послідовного навчання слабких моделей. Кожне нове дерево намагається виправити помилки попередніх, використовуючи градієнтний спуск для оновлення ваг. Однак цей підхід чутливий до масштабу ознак, адже градієнтний спуск враховує відносні зміни значень, що може спотворювати процес навчання.

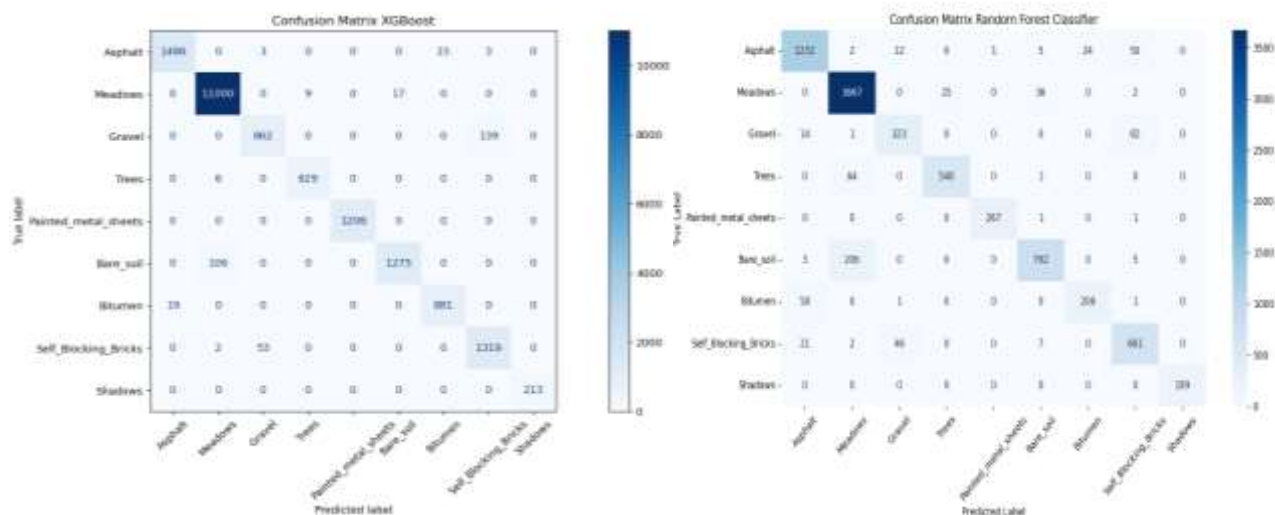


Рис. 3. Матриці невідповідностей алгоритмів XGBoost та *RandomForestClassifier*

На відміну від нього, *RandomForestClassifier* застосовує метод бэггінгу (bagging), у якому кожне дерево будується незалежно на випадковій підвибірці даних. Завдяки цьому модель краще справляється з неоднорідними масштабами ознак і є менш чутливою до викидів. Аналізуючи матриці невідповідностей (рис. 3), можна побачити, що *RandomForestClassifier* точніше класифікує дрібні категорії матеріалів, тоді як XGBoost часто помилково відносить інші матеріали до класу "Meadows". Однак, якщо розглядати матрицю відповідностей, XGBoost демонструє вищу продуктивність у більшості класів. Виникає суперечливість між точністю класифікації та загальними характеристиками моделей.

З метою отримання об'єктивнішої оцінки було проаналізовано додаткові метрики: *precision* (точність), *recall* (повнота) та *f1-score* (середнє значення точності та повноти). Вони дозволяють краще оцінити баланс між правильними позитивними та негативними класифікаціями (табл. 1). *RandomForestClassifier* показує кращу стійкість до шуму і змішаних класів, що робить його більш придатним для аналізу гіперспектральних зображень, де дані часто мають складну структуру та значну варіативність.

Результати класифікації вхідних даних за методами
XGBoost та Random Forest Classifier

<i>Назва методу</i>	<i>Random Forest Classifier</i>			<i>XGBoost</i>		
<i>матеріали / метрики</i>	<i>precision</i>	<i>recall</i>	<i>f1-score</i>	<i>precision</i>	<i>recall</i>	<i>f1-score</i>
<i>Asphalt</i>	0.93	0.93	0.93	0.99	0.24	0.38
<i>Meadows</i>	0.93	0.98	0.96	0.99	0.59	0.74
<i>Gravel</i>	0.85	0.77	0.81	0.94	0.4	0.57
<i>Trees</i>	0.96	0.89	0.92	1	0.21	0.34
<i>Painted metal sheets</i>	1	0.99	0.99	1	0.9	0.95
<i>Bare Soil</i>	0.94	0.79	0.86	0.98	0.25	0.39
<i>Bitumen</i>	0.9	0.77	0.83	0.96	0.64	0.77
<i>Self Blocking Bricks</i>	0.82	0.9	0.86	0.92	0.36	0.52
<i>Shadows</i>	1	1	1	1	0.22	0.36

*Джерело розроблене авторами

Аналізуючи показники впевненості моделей у власних прогнозах (precision) та здатності правильно ідентифікувати об'єкти (recall), можна зробити висновок, що RandomForestClassifier точніше визначає свої передбачення. Ця модель демонструє високі значення precision, recall та f1-score для всіх класів, що вказує на її збалансованість та стабільність у розпізнаванні різних матеріалів.

На противагу, XGBoost має значно нижчі значення recall та f1-score, що свідчить про більшу кількість неправильно класифікованих об'єктів. Наприклад, при розпізнаванні асфальту f1-score для RandomForestClassifier становить 0.93, тоді як для XGBoost – лише 0.38. Попри високий precision (0.99), у XGBoost дуже низький recall (0.24), що вказує на те, що модель часто не розпізнає цей клас, тобто багато пікселів, що належать до асфальту, залишаються некласифікованими або віднесеними до інших класів. Аналогічна ситуація спостерігається і для інших категорій, де XGBoost демонструє низькі recall та f1-score, що вказує на значні пропуски у передбаченнях.

Незважаючи на те, що матриця невідповідностей може створювати враження вищої точності XGBoost, детальніший аналіз додаткових метрик, зокрема precision, recall та f1-score, показує протилежне. RandomForestClassifier забезпечує більш точну і стабільну класифікацію, оскільки не лише правильно класифікує більшість пікселів, але й демонструє збалансоване співвідношення між precision та recall. Це робить його більш ефективним для задачі класифікації гіперспектральних зображень, особливо в умовах нерівномірного розподілу даних та присутності шуму.

В результаті, наведена модель системи ідентифікації матеріалів на земній поверхні шляхом класифікації гіперспектральних зображень. Програма, розроблена в рамках дослідження, здійснює зменшення розмірності вхідних даних, їх нормалізацію, тренування моделі за допомогою крос-валідації та подальшу класифікацію. Вихідними результатами є прогнозоване зображення, графіки, що візуалізують розпізнані матеріали, їх кількість і спектральні характеристики.

Основна ідея полягає у запровадженні інноваційних підходів до класифікації матеріалів на основі гіперспектрального аналізу, у поєднанні методів зменшення розмірності з алгоритмами машинного навчання. Це дослідження розкриває як теоретичні аспекти, так і практичну значущість застосування гіперспектральних методів, демонструючи їх потенціал для підвищення точності розпізнавання об'єктів у різних сферах, зокрема для військової розвідки, екологічного моніторингу та медичної діагностики. Отже, розроблена система класифікації дозволяє оптимізувати обчислювальні ресурси та забезпечити високоточне визначення матеріалів, що є основою для подальших досліджень і практичного впровадження цієї технології.

Висновки і перспективи подальших досліджень. Методи класифікації, розглянуті в даному дослідженні, базуються на аналізі спектральних характеристик кожного пікселя, що дозволяє точно ідентифікувати матеріали на гіперспектральних зображеннях. Порівняльний аналіз алгоритмів машинного навчання показав, що *Random Forest Classifier* демонструє вищу ефективність у класифікації об'єктів порівняно з *XGBoost*. Це зумовлено його здатністю краще працювати з неоднорідними даними, знижувати похибки класифікації та підвищувати стійкість до надмірного навчання.

На основі отриманих результатів розробляється веб-застосунок для автоматизованої обробки гіперспектральних зображень. Його *backend* реалізовано за допомогою *Flask* – фреймворку на *Python*, який дозволяє ефективно обробляти вхідні зображення, зберігати результати класифікації та взаємодіяти з базою даних. *Frontend* розроблено з використанням *HTML*, *CSS* та *JavaScript* забезпечує зручний інтерфейс для візуалізації результатів аналізу, перегляду класифікованих зображень та їхньої аналітики.

Отримані результати демонструють високу ефективність застосування методу *Random Forest Classifier* для класифікації матеріалів на гіперспектральних зображеннях. Запропонований підхід дозволяє підвищити точність розпізнавання та скоротити обчислювальні витрати. Розроблений веб-застосунок сприяє автоматизації процесу аналізу, що значно розширює можливості використання гіперспектральних даних у різних сферах.

Подальші дослідження зосереджені на підвищенні обчислювальної ефективності алгоритмів для обробки великих обсягів гіперспектральних даних у реальному часі. Особлива увага приділяється інтеграції глибоких нейронних мереж, що дозволить покращити точність класифікації та адаптацію моделей до різних типів матеріалів. Передбачається розширення функціоналу веб-застосунку шляхом впровадження додаткових аналітичних інструментів і підтримки різних форматів гіперспектральних даних. Окрім того, вивчаються можливості застосування розробленої технології у сферах військової розвідки, екологічного моніторингу, дистанційного зондування та пошуку небезпечних об'єктів, таких як міни чи вибухонебезпечні предмети. Наукова новизна дослідження полягає у поєднанні передових методів машинного навчання з автоматизованою обробкою гіперспектральних даних, що розширює можливості їх практичного застосування.

Список використаної літератури

1. Amigo J. M. Hyperspectral Imaging (Data Handling in Science and Technology, Volume 32). – 1st Edition [Електронний ресурс]. – Режим доступу: <https://tinyurl.com/4v2c37sj>. – Дата публікації: 15.09.2019.
2. FlySight. Understanding Hyperspectral Imaging - Benefits, Use Cases And Examples [Електронний ресурс]. – Режим доступу: <https://www.flysight.it/understanding-hyperspectral-imaging-benefits-use-cases-examples/>. – Дата публікації: 23.08.2023.
3. Jaadi Z. Principal Component Analysis (PCA): A Step-by-Step Explanation [Електронний ресурс]. – Режим доступу: <https://builtin.com/data-science/step-step-explanation-principal-component-analysis>. – Дата звернення: 29.01.2025.
4. Jaiswal G., Rani R., Mangotra H., Sharma A. Integration of hyperspectral imaging and autoencoders: Benefits, applications, hyperparameter tuning and challenges // Bennett University, Greater Noida, India, Indira Gandhi Delhi Technical University for Women, Delhi, India [Електронний ресурс]. – Режим доступу: <https://www.sciencedirect.com/science/article/abs/pii/S1574013723000515>. – Дата публікації: 31.08.2023.
5. Kaggle. Pavia University HSI [Електронний ресурс]. – Режим доступу: <https://tinyurl.com/2yz68pav>. – Дата звернення: 29.01.2025.
6. Lone Z. A., Pais A. R. Object detection in hyperspectral images [Електронний ресурс]. – Режим доступу: <https://www.sciencedirect.com/science/article/abs/pii/S1051200422003694>. – Дата публікації: 05.10.2022.

7. Medium. Making Data Simpler with Principal Component Analysis (PCA) [Електронний ресурс]. – Режим доступу: <https://medium.com/@deepyachowdary/making-data-simpler-with-principal-component-analysis-pca-20a910a449d6>. – Дата публікації: 25.01.2024.
8. V5 Geospatial Solutions, Inc. Basic Hyperspectral Analysis Tutorial [Електронний ресурс]. – Режим доступу: <https://www.nv5geospatialsoftware.com/docs/HyperspectralAnalysisTutorial.html>. – Дата звернення: 02.01.2025.
9. XGBoost. XGBoost Documentation. [Електронний ресурс]. – Режим доступу: <https://xgboost.readthedocs.io/en/stable/>.
10. Чорнобривець Д. В., Поперешняк С. В., Каплюк В. О. Система моніторингу місць для паркування з використанням комп'ютерного зору [Електронний ресурс] // Науковий журнал «Телекомунікації та інформаційні технології». 2024, №4 (85). С. 72-82
11. Алексіна, Л. Т., & Бондарчук, А. П. (2024). Оптимізація гіперпараметрів для машинного навчання. Зв'язок, (2), 18-22.
<https://con.dut.edu.ua/index.php/communication/article/view/2755>
12. Бондарчук, А. П., Онисько, А. І., Отрох, С. І., & Шевчук, Д. О. (2023). Система двофакторної аутентифікації користувача за допомогою розпізнавання обличчя. Телекомунікаційні та інформаційні технології, (3), 79-84.
<https://tit.dut.edu.ua/index.php/telecommunication/article/view/2481/2363>