

**МІНІСТЕРСТВО ВНУТРІШНІХ СПРАВ УКРАЇНИ
ЛЬВІВСЬКИЙ ДЕРЖАВНИЙ УНІВЕРСИТЕТ ВНУТРІШНІХ СПРАВ
НАВЧАЛЬНО-НАУКОВИЙ ІНСТИТУТ УПРАВЛІННЯ,
ПСИХОЛОГІЇ ТА БЕЗПЕКИ
Кафедра інформаційних технологій**

**Моделі машинного навчання для прогнозування дій
супротивника в умовах неповної та невизначеної інформації**

кваліфікаційна робота

здобувача вищої освіти

4 курсу денної форми навчання

Данила Богатирьова

Науковий керівник:

доктор філософії

Олег БАСИСТЮК

Рецензент:

Кваліфікаційна робота допущена до захисту

«___» _____ 2026 р., протокол № _____

Завідувач кафедри інформаційних технологій

_____ **Олег ЗАЧЕК**

АНОТАЦІЯ

Кваліфікаційна робота присвячена розробці системи прогнозування дій супротивника на основі моделей машинного навчання в умовах неповної та невизначеної інформації.

У роботі проведено аналіз сучасних методів прогнозування поведінки складних систем, досліджено особливості застосування моделей машинного навчання та виконано аналіз методів роботи з неповними даними.

У межах роботи розроблено програмний модуль прогнозування, який реалізує автоматичну генерацію синтетичного набору даних, механізми моделювання невизначеності, попередню обробку інформації, навчання моделей машинного навчання та оцінювання результатів.

Для реалізації системи використано мову програмування Python та бібліотеки pandas, NumPy, scikit-learn, matplotlib, PyQt5.

Проведено експериментальне дослідження роботи системи, виконано порівняння моделей машинного навчання та визначено найбільш ефективний алгоритм прогнозування.

Результати дослідження показали ефективність використання моделей машинного навчання для прогнозування дій супротивника навіть за наявності неповної інформації.

Ключові слова: машинне навчання, прогнозування, невизначеність, штучний інтелект, класифікація, Python, Random Forest, аналіз даних.

ABSTRACT

The qualification work is devoted to the development of a system for predicting enemy actions using machine learning models under conditions of incomplete and uncertain information.

The paper analyzes modern approaches to behavior prediction in complex systems and investigates the application of machine learning methods for classification tasks. Methods for handling incomplete and uncertain data were also considered.

Within the practical part of the work, a software module was developed that implements automatic generation of a synthetic dataset, uncertainty simulation mechanisms, data preprocessing, model training and evaluation of prediction results.

The system was implemented using Python programming language and the following libraries: pandas, NumPy, scikit-learn, matplotlib and PyQt5.

Experimental studies were conducted to compare machine learning models and determine the most effective prediction algorithm.

The obtained results demonstrated the effectiveness of machine learning models for predicting enemy actions even under incomplete information conditions.

Keywords: machine learning, prediction, uncertainty, artificial intelligence, classification, Python, Random Forest, data analysis.

Зміст

ПЕРЕЛІК УМОВНИХ СКОРОЧЕНЬ.....	5
ВСТУП.....	6
РОЗДІЛ 1. АНАЛІЗ МЕТОДІВ ПРОГНОЗУВАННЯ ДІЙ СУПРОТИВНИКА В УМОВАХ НЕПОВНОЇ ТА НЕВИЗНАЧЕНОЇ ІНФОРМАЦІЇ.....	8
1.1 Аналіз предметної області прогнозування дій супротивника.....	8
1.2 Аналіз моделей машинного навчання для задач прогнозування.....	10
1.3 Методи роботи з неповними та невизначеними даними.....	14
РОЗДІЛ 2. ПРОЄКТУВАННЯ ПРОГРАМНОЇ СИСТЕМИ ПРОГНОЗУВАННЯ...	17
2.1 Обґрунтування вибору засобів реалізації.....	17
2.2 Проєктування архітектури програмної системи.....	18
2.3 Проєктування структури синтетичного набору даних.....	21
2.4 Проєктування механізму моделювання невизначеності.....	23
2.5 Проєктування алгоритму попередньої обробки даних.....	26
2.6 Проєктування алгоритму функціонування системи.....	29
2.7 Проєктування інтерфейсу користувача.....	32
Висновки до розділу 2.....	36
РОЗДІЛ 3. РЕАЛІЗАЦІЯ ПРОГРАМНОГО МОДУЛЯ ТА ЕКСПЕРИМЕНТАЛЬНЕ ДОСЛІДЖЕННЯ.....	37
3.1 Реалізація генерації синтетичного набору даних.....	37
3.2 Реалізація механізму пошкодження даних та моделювання невизначеності....	39
3.3 Реалізація попередньої обробки даних.....	42

3.4 Реалізація моделей машинного навчання.....	46
3.5 Навчання та тестування моделей.....	51
3.6 Реалізація графічного інтерфейсу та аналіз результатів прогнозування.....	54
3.7 Аналіз результатів експериментального дослідження.....	59
Висновки до розділу 3.....	60
ВИСНОВКИ.....	61
СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ.....	63
ДОДАТОК.....	65

ПЕРЕЛІК УМОВНИХ СКОРОЧЕНЬ

ML — Machine Learning (машинне навчання)

AI — Artificial Intelligence (штучний інтелект)

GUI — Graphical User Interface (графічний інтерфейс користувача)

CSV — Comma-Separated Values

NaN — Not a Number

API — Application Programming Interface

RF — Random Forest

DT — Decision Tree

LR — Logistic Regression

NB — Naive Bayes

UI — User Interface

ВСТУП

Сучасний розвиток інформаційних технологій та систем штучного інтелекту створює нові можливості для автоматизації процесів аналізу та прогнозування складних подій у різних сферах діяльності. Особливої актуальності набуває застосування методів машинного навчання для підтримки прийняття рішень в умовах швидкої зміни ситуації, значного обсягу інформації та наявності факторів невизначеності. Однією з важливих задач є прогнозування можливих дій супротивника на основі наявних даних, які часто характеризуються неповнотою, неточністю або наявністю інформаційного шуму. У реальних умовах повна інформація про наміри, ресурси та поведінку супротивника практично недоступна, що значно ускладнює процес аналізу та прийняття рішень.

Проблематика прогнозування поведінки супротивника розглядалася багатьма українськими та зарубіжними науковцями у сфері штучного інтелекту, аналізу даних та систем підтримки прийняття рішень. Значний внесок у розвиток методів машинного навчання здійснили Arthur Samuel, Tom Mitchell, Geoffrey Hinton та інші. Дослідження сучасних алгоритмів класифікації, прогнозування та обробки невизначених даних демонструють перспективність використання ансамблевих моделей, статистичних методів та методів інтелектуального аналізу даних.

Метою роботи є розроблення програмного модуля прогнозування дій супротивника на основі моделей машинного навчання в умовах неповної та невизначеної інформації.

Для досягнення поставленої мети необхідно виконати такі завдання:

- дослідити сучасні підходи до прогнозування поведінки об'єктів у складних інформаційних системах;
- проаналізувати існуючі методи машинного навчання для задач класифікації;
- дослідити особливості роботи з неповними та невизначеними даними;
- сформулювати набір параметрів для моделювання дій супротивника;
- розробити структуру навчального набору даних;
- реалізувати алгоритми попередньої обробки інформації;
- реалізувати програмний модуль прогнозування;
- виконати навчання моделей машинного навчання;
- провести порівняльний аналіз результатів;
- визначити найбільш ефективний метод прогнозування.

Об'єктом дослідження є процес прогнозування поведінки супротивника в умовах неповної інформації.

Предметом дослідження є моделі машинного навчання для визначення та прогнозування ймовірних дій супротивника.

Під час виконання роботи використовувалися методи системного аналізу, методи статистичного аналізу даних, алгоритми машинного навчання, методи класифікації, методи математичного моделювання та програмні засоби обробки даних мовою Python.

Практичне значення роботи полягає у створенні програмного модуля, який дозволяє прогнозувати можливі дії супротивника на основі неповних вхідних даних та може бути використаний як елемент систем підтримки прийняття рішень.

Кваліфікаційна робота складається зі вступу, трьох розділів, висновків, списку використаних джерел та додатків.

РОЗДІЛ 1 АНАЛІЗ СУЧАСНИХ МЕТОДІВ ПРОГНОЗУВАННЯ ДІЙ СУПРОТИВНИКА В УМОВАХ НЕВИЗНАЧЕНОСТІ

1.1 Аналіз задачі прогнозування поведінки супротивника

Прогнозування поведінки складних об'єктів є однією з найбільш важливих задач інформаційно-аналітичних систем. Основною метою прогнозування є визначення ймовірних майбутніх дій об'єкта на основі наявних даних про його попередню поведінку та поточний стан середовища.

Процес прогнозування ускладнюється тим, що інформація про об'єкт спостереження часто має неповний характер. Частина параметрів може бути недоступною, втраченою або містити помилки. Крім того, інформація може надходити із різних джерел, що відрізняються ступенем достовірності.

До основних факторів невизначеності належать:

- неповнота інформації;
- наявність шуму у вхідних даних;
- відсутність частини параметрів;
- часові затримки надходження інформації;
- випадковість поведінки об'єкта;
- зовнішні фактори впливу.

У задачах прогнозування дій супротивника можуть використовуватись такі параметри:

- відстань до об'єкта;
- кількість сил супротивника;
- тип місцевості;
- погодні умови;
- видимість;
- рівень ризику;
- попередні дії;
- час доби;
- стан зв'язку;
- втрати.

На основі аналізу зазначених параметрів формується набір ознак, який використовується для навчання моделей машинного навчання.

Для даної роботи передбачається створення синтетичного набору даних, що моделюватиме типові ситуації та дозволить виконати навчання моделей без використання реальних спеціалізованих даних.



1.2 Аналіз моделей машинного навчання для задач прогнозування

Одним із найбільш перспективних напрямків розвитку сучасних інформаційних технологій є використання моделей машинного навчання для аналізу великих масивів даних та прогнозування подальшого розвитку подій. Методи машинного навчання широко застосовуються у системах підтримки прийняття рішень, інтелектуальному аналізі інформації, кібербезпеці, прогнозуванні поведінки складних систем та інших галузях.

Основною перевагою машинного навчання є можливість виявлення прихованих закономірностей між параметрами системи без побудови жорстко визначених алгоритмів. Замість використання наперед заданих правил модель формує залежності між вхідними параметрами та очікуваним результатом на основі аналізу навчальних даних.

Для задач прогнозування поведінки супротивника в умовах неповної та невизначеної інформації найбільш доцільним є використання алгоритмів класифікації. У межах даної роботи розглядаються такі моделі:

- дерево рішень (Decision Tree);
- випадковий ліс (Random Forest);
- логістична регресія (Logistic Regression);
- наївний байєсівський класифікатор (Naive Bayes).

1.2.1 Метод дерева рішень

Дерево рішень є одним із найбільш поширених алгоритмів класифікації. Його принцип роботи полягає у послідовному розбитті набору даних на окремі групи відповідно до певних умов.

Модель представляється у вигляді ієрархічної структури, де:

- вершини містять перевірку певної ознаки;
- гілки описують можливі результати;
- листки містять кінцевий прогноз.

Наприклад, система може послідовно перевіряти:

1. рівень видимості;
2. кількість супротивників;
3. попередні дії;
4. відстань.

Після проходження всіх умов модель формує прогнозовану дію.

Перевагами дерева рішень є	До недоліків належать
простота реалізації	схильність до перенавчання
висока швидкість роботи	залежність результатів від структури навчальних даних
наочність	нестабільність при появі шуму
можливість візуалізації результатів	

Для задач із неповними даними використання одного дерева рішень може давати недостатню точність.

1.2.2 Метод випадкового лісу

Метод випадкового лісу є ансамблевим алгоритмом машинного навчання, що базується на використанні множини дерев рішень.

Принцип роботи Random Forest полягає у побудові декількох незалежних дерев рішень на випадкових підмножинах даних. Кінцевий результат визначається методом голосування між усіма побудованими моделями.

Загальний алгоритм роботи має вигляд:

1. створення випадкової вибірки;
2. побудова дерева рішень;
3. повторення процедури багаторазово;
4. обчислення результату голосування.

Переваги Random Forest	Недоліки
висока точність прогнозування	більша складність
стійкість до шуму	потреба у більшій кількості ресурсів
робота з великою кількістю ознак	зниження інтерпретованості
менша ймовірність перенавчання	
можливість оцінювання важливості ознак	

У задачах з неповною інформацією Random Forest демонструє кращі результати порівняно зі звичайними деревами рішень.

1.2.3 Логістична регресія

Логістична регресія належить до статистичних методів класифікації та використовується для оцінювання ймовірності належності об'єкта до певного класу.

Алгоритм дозволяє визначити залежність між вхідними параметрами та ймовірністю появи конкретної події.

Оснoву моделі становить логістична функція:

$$P(y) = \frac{1}{1 + e^{-z}} \quad P(y) = \frac{1}{1 + e^{-z}}$$

де:

$P(y)$ — ймовірність події;

z — сукупність вхідних параметрів.

Переваги	Недоліки
простота реалізації	складність роботи зі складними нелінійними залежностями
висока швидкість навчання	зниження точності при великій кількості шумових даних
можливість інтерпретації результатів	
низькі вимоги до ресурсів	

1.2.4 Наївний байєсівський класифікатор

Наївний байєсівський класифікатор належить до групи ймовірнісних моделей машинного навчання. Алгоритм базується на застосуванні теореми Байєса та припущенні про незалежність ознак.

Основою методу є:

$$P(A|B) = \frac{P(B|A)P(A)}{P(B)}$$

де:

$P(A|B)$ — ймовірність події A за умови B ;

$P(B|A)$ — умовна ймовірність;

$P(A)$ — апіорна ймовірність;

$P(B)$ — загальна ймовірність.

Переваги	Недоліки
швидке навчання	припущення про незалежність параметрів
ефективність при невеликих наборах даних	зниження точності для складних взаємозалежних даних
стійкість до пропусків	

Проведений аналіз показав, що кожна модель має власні переваги та недоліки. Для подальшого дослідження доцільно реалізувати декілька моделей та провести експериментальне порівняння їх ефективності в умовах неповної та невизначеної інформації. Очікується, що ансамблевий алгоритм Random Forest забезпечить найвищу точність прогнозування.

1.3 Аналіз методів роботи з неповними та невизначеними даними

Однією з основних проблем під час створення систем прогнозування є робота з неповними даними. У реальних умовах інформація практично ніколи не надходить у повному вигляді. Частина параметрів може бути втрачена, отримана із затримкою або містити неточності. Через це система повинна не лише аналізувати інформацію, а й уміти працювати в умовах невизначеності.

Під невизначеністю будемо розуміти ситуацію, коли частина характеристик об'єкта невідома або має недостатній рівень достовірності. Наприклад, може бути невідома точна кількість супротивників, якість зв'язку або реальна дистанція до цілі. Крім того, одна й та сама інформація може надходити з декількох джерел і відрізнятися між собою.

На практиці проблеми вхідних даних найчастіше поділяють на декілька категорій:

- відсутні значення;
- пошкоджені дані;
- шумові значення;
- дублювання інформації;
- суперечливі дані;
- випадкові помилки.

Відсутність даних є однією з найпоширеніших проблем. У наборах даних вона зазвичай позначається спеціальними значеннями Null або NaN. Такі ситуації можуть виникати через несправності обладнання, втрату інформації під час передачі або людський фактор.

Для роботи з пропущеними значеннями використовуються декілька підходів.

Перший метод полягає у видаленні записів, які містять пропуски. Такий підхід є простим, але може призвести до втрати значного обсягу інформації.

Другий підхід передбачає заповнення пропусків певними значеннями. Наприклад:

- середнім значенням;
- медіаною;

- найбільш поширеним значенням;
- результатом прогнозування.

У роботі планується використати метод заповнення пропусків через статистичні характеристики вибірки.

Окрему проблему становить шум у даних. Шумом називають випадкові спотворення інформації, які можуть впливати на результати навчання моделі.

Наприклад, фактична кількість супротивників може становити:

10 осіб,

але система отримала:

12 осіб або 8 осіб.

На перший погляд різниця невелика, проте при великому обсязі даних навіть незначні відхилення можуть суттєво впливати на результати прогнозування.

Для моделювання умов невизначеності в межах цієї роботи планується спеціально додавати шум у навчальний набір даних. Це дозволить перевірити, наскільки стійкими є моделі машинного навчання до викривленої інформації.

Під час розроблення програмного забезпечення також передбачається штучне створення неповних даних шляхом генерації пропущених значень:

```
data.loc[random_rows,"visibility"]=np.nan  
data.loc[random_rows,"enemy_count"]=np.nan
```

Такий підхід дозволить наблизити модель до реальних умов функціонування.

Після заповнення пропусків та усунення частини шуму дані проходитимуть етап попередньої обробки. До нього входитимуть:

- перетворення категоріальних ознак;
- нормалізація числових параметрів;
- масштабування;
- підготовка навчальної та тестової вибірки.

Таким чином, якість роботи моделей машинного навчання значною мірою залежить не лише від вибраного алгоритму, а й від правильності підготовки даних. Навіть хороша модель може показувати слабкі результати при роботі з необробленими або пошкодженими даними.

РОЗДІЛ 2

ВИБІР ЗАСОБІВ ТА ПРОЄКТУВАННЯ ПРОГРАМНОЇ СИСТЕМИ

2.1 Обґрунтування вибору мови програмування та середовища розробки

Для реалізації програмної системи прогнозування дій супротивника було обрано мову програмування Python. Вибір саме цієї мови зумовлений її широким використанням у сфері аналізу даних, штучного інтелекту та машинного навчання.

Сьогодні Python фактично став одним із основних інструментів для створення систем машинного навчання. Причина доволі проста: мова має велику кількість готових бібліотек, що дозволяють швидко створювати моделі, обробляти дані та проводити експерименти без необхідності реалізації складних алгоритмів з нуля.

Ще однією перевагою Python є зрозумілий синтаксис. У порівнянні з багатьма іншими мовами програмування код виглядає компактнішим і легше читається. Це спрощує подальшу підтримку програмного забезпечення та внесення змін у проєкт.

Під час реалізації системи використовувалися такі можливості Python:

- створення структур даних;
- реалізація алгоритмів машинного навчання;
- побудова графіків;
- створення графічного інтерфейсу;
- збереження результатів.

Розроблена система складається з двох основних частин:

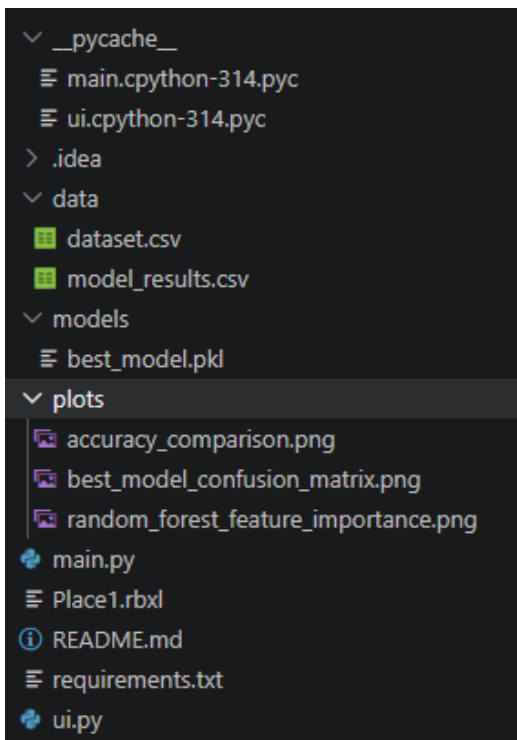
- програмного модуля обробки та навчання моделей;
- графічного інтерфейсу користувача.

Основна логіка системи реалізована у файлі `main.py`, який відповідає за створення датасету, навчання моделей та формування результатів.

Графічна взаємодія користувача із системою реалізована у файлі `ui.py`.

Для розробки використовувалося сучасне середовище програмування. Воно забезпечує:

- підсвічування синтаксису;
- автоматичне форматування коду;
- налагодження;
- підтримку бібліотек;
- інтеграцію з пакетним менеджером Python.



У результаті вибір Python дозволив скоротити час розробки та зосередитися безпосередньо на реалізації алгоритмів прогнозування.

2.2 Вибір бібліотек для реалізації системи

Під час розроблення програмного забезпечення важливим етапом є вибір бібліотек та програмних компонентів, які забезпечують реалізацію необхідного функціоналу системи. Оскільки основною задачею роботи є прогнозування дій супротивника з використанням методів машинного навчання, виникла необхідність використання інструментів для обробки даних, навчання моделей, візуалізації результатів та створення графічного інтерфейсу.

Для реалізації системи було використано декілька бібліотек мови Python. Кожна з них виконує окрему функцію та відповідає за конкретний етап роботи програми.

Однією з основних бібліотек є **pandas**, яка використовується для створення та обробки таблиць даних. У програмі бібліотека застосовується для формування синтетичного датасету, зчитування інформації та подальшого збереження результатів у формат CSV.

Бібліотека **NumPy** використовується для виконання математичних операцій та роботи з багатовимірними масивами. Крім цього, вона застосовується для генерації випадкових чисел, створення шуму та моделювання невизначеності в наборі даних.

Для реалізації моделей машинного навчання використовується бібліотека **scikit-learn**. Вона містить готові інструменти для:

- навчання моделей;
- попередньої обробки даних;
- поділу вибірки;
- оцінювання результатів;
- побудови конвеєрів обробки.

У розробленій системі застосовуються такі моделі:

- Decision Tree;
- Random Forest;
- Logistic Regression;
- Gaussian Naive Bayes.

Також бібліотека використовується для реалізації механізму попередньої обробки даних. Для цього застосовуються компоненти:

- SimpleImputer;
- StandardScaler;
- OneHotEncoder;
- Pipeline;
- ColumnTransformer.

Якщо говорити простими словами, то ця частина працює як автоматичний конвеєр. Дані не передаються напряму в модель. Спочатку система знаходить пропуски, потім заповнює їх, переводить текстові значення у числовий вигляд, масштабує параметри і лише після цього запускає навчання.

Для побудови графіків використовуються бібліотеки **matplotlib** та **seaborn**. Вони дозволяють створювати:

- графік порівняння точності моделей;
- матрицю помилок;
- графік важливості ознак.

Результати автоматично зберігаються у директорію графічних матеріалів та можуть використовуватись для подальшого аналізу.

Для створення графічного інтерфейсу користувача використовується бібліотека **PyQt5**. Вона забезпечує створення сучасного віконного інтерфейсу із підтримкою кнопок, полів введення, графічних компонентів та системи подій.

У програмі реалізовано:

- панель введення параметрів;
- блок результатів;
- область візуалізації;
- журнал подій;
- механізм збереження результатів.

Крім цього, у проєкті використовується бібліотека **joblib**, яка дозволяє зберігати навчену модель у файл для подальшого використання без повторного навчання.

Це дозволяє значно скоротити час запуску системи та уникнути повторних обчислень.

Отже, обраний набір бібліотек повністю забезпечує реалізацію функціональних можливостей системи та дозволяє виконати всі поставлені завдання дипломної роботи.

2.3 Проєктування архітектури системи

На етапі проєктування програмної системи важливо визначити структуру її компонентів та принцип взаємодії між окремими модулями. Правильно побудована архітектура спрощує розробку, подальшу підтримку програмного забезпечення та внесення змін до окремих частин системи.

У межах даної роботи реалізовано модульну архітектуру. Такий підхід дозволяє розділити функціональні частини програми на незалежні компоненти, кожен із яких виконує власні завдання.

Розроблена система складається з таких основних частин:

- модуль генерації даних;
- модуль попередньої обробки;
- модуль навчання моделей;
- модуль оцінювання результатів;
- модуль збереження даних;
- графічний інтерфейс користувача.

Загальна структура програмної системи показана у вигляді взаємодії між двома основними файлами:

- main.py;
- ui.py.

Файл main.py реалізує основну логіку роботи системи. Саме тут відбувається:

- створення синтетичного набору даних;
- генерація невизначеності;
- навчання моделей;
- побудова графіків;
- оцінювання результатів;
- збереження моделей.

У файлі ui.py реалізовано графічний інтерфейс користувача, який взаємодіє з функціями бекенду.

Якщо описати це простіше, то main.py є "мозком" системи, а ui.py — оболонкою, через яку користувач взаємодіє з програмою.

Користувач не працює безпосередньо з алгоритмами машинного навчання. Він вводить параметри через інтерфейс, а система самостійно виконує всі необхідні дії.



Послідовність роботи виглядає таким чином:

1. користувач вводить параметри;
2. система формує набір вхідних даних;
3. виконується попередня обробка;
4. дані передаються у модель;
5. формується прогноз;
6. результат відображається на екрані.

Окремим елементом архітектури є механізм збереження результатів роботи системи. Після завершення навчання автоматично створюються:

- файл набору даних;
- файл результатів моделей;
- файл збереженої моделі;
- графічні матеріали.

У коді для цього використовуються окремі директорії: /data,/models,

/plots

Використання такого підходу дозволяє відокремити результати роботи від програмного коду та спрощує подальше використання системи.

2.4 Формування структури синтетичного набору даних

Одним із ключових етапів побудови моделей машинного навчання є формування набору даних. Від якості вхідної інформації значною мірою залежить точність та ефективність подальшого прогнозування.

Для навчання моделей у межах даної роботи використовується синтетичний набір даних. Вибір такого підходу зумовлений тим, що використання реальних спеціалізованих даних може бути обмеженим або недоступним. Крім того, створення штучного набору дозволяє контролювати характеристики даних та моделювати різні ситуації.

У створеній системі датасет генерується автоматично. Формування даних реалізовано через окрему функцію генерації.

Програмно генерується набір із 1200 записів, кожен із яких описує окрему умовну тактичну ситуацію.

Кожний запис містить набір параметрів, які використовуються моделлю для прогнозування подальших дій.

Структура набору даних складається з таких ознак:

Назва параметра	Опис
distance_to_enemy	відстань до супротивника
enemy_count	кількість супротивників
terrain_type	тип місцевості
visibility_level	рівень видимості
ammo_level	рівень боєзапасу
previous_enemy_action	попередня дія
time_of_day	час доби
risk_level	рівень ризику
losses_detected	рівень втрат
communication_quality	якість зв'язку
next_enemy_action	цільова змінна

Цільова змінна `next_enemy_action` використовується для навчання моделей машинного навчання та визначає прогнозовану наступну дію супротивника.

Можливі значення:

- attack;
- retreat;
- reconnaissance;
- flanking;
- waiting.

Якщо пояснювати простіше, то система отримує набір характеристик ситуації та намагається визначити, що супротивник зробить далі.

Наприклад:

відстань — 25 км;
видимість — низька;
кількість супротивників — 12;
ризик — низький;
боєзапас — високий.

На основі цих параметрів система може зробити висновок, що найбільш імовірною буде дія «attack».

Під час генерації даних використовується набір умовних залежностей між параметрами та результатом прогнозування. Для цього реалізовано механізм розрахунку ймовірностей.

Наприклад:

при невеликій відстані та високому боєзапасі збільшується ймовірність атаки;
при великих втратах підвищується ймовірність відступу;
низька видимість збільшує шанс проведення розвідки.

Таким чином система формує більш реалістичні статистичні залежності.

Після генерації основних параметрів у системі виконується додаткове моделювання невизначеності. Частина даних спеціально пошкоджується шляхом:

- створення пропущених значень;
- додавання випадкового шуму;
- формування невизначених категорій.

Це дозволяє імітувати ситуації, максимально наближені до реальних умов функціонування систем аналізу інформації.

Завдяки використанню синтетичного набору даних стало можливим створити контрольоване середовище для проведення експериментів та подальшого аналізу роботи моделей машинного навчання.

2.5 Проєктування алгоритму попередньої обробки даних

Перед навчанням моделей машинного навчання дані необхідно підготувати. Це важливий етап, оскільки більшість алгоритмів не можуть якісно працювати з пропущеними значеннями, текстовими параметрами або ознаками, які мають різні масштаби.

У розробленій системі попередня обробка даних виконується автоматично за допомогою спеціального конвеєра. Такий підхід дозволяє не обробляти кожну ознаку вручну, а створити загальну схему підготовки даних.

У програмі всі ознаки поділено на дві групи:

- числові ознаки;
- категоріальні ознаки.

До числових ознак належать:

- distance_to_enemy;
- enemy_count;
- visibility_level;
- ammo_level;
- risk_level;
- losses_detected;
- communication_quality.

До категоріальних ознак належать:

- terrain_type;
- previous_enemy_action;
- time_of_day.

Для числових ознак використовується заповнення пропущених значень медіаною. Це означає, що якщо певне числове значення відсутнє, система автоматично замінює його центральним значенням вибірки.

Такий підхід є зручним, бо медіана менш чутлива до різких викидів, ніж середнє арифметичне.

Після заповнення пропусків числові ознаки масштабуються за допомогою StandardScaler. Це потрібно для того, щоб різні параметри мали приблизно однаковий вплив на модель.

Наприклад, відстань може вимірюватися у межах від 0 до 150, а рівень ризику — від 0 до 100. Без масштабування модель може сприймати деякі ознаки як більш важливі лише через більші числові значення.

Для категоріальних ознак використовується інший підхід. Спочатку пропуски заповнюються найбільш поширеним значенням, а потім текстові категорії перетворюються у числовий формат за допомогою OneHotEncoder.

Це потрібно тому, що моделі машинного навчання не працюють напряму з текстом. Наприклад, значення:

Urban, forest, mountain

перетворюються у набір числових колонок.

Якщо простіше, програма переводить текст у формат, який може зрозуміти модель.

У коді для цього використовується ColumnTransformer, який дозволяє застосувати різні способи обробки до різних типів ознак.

Загальний алгоритм попередньої обробки має такий вигляд:

1. отримання вхідного набору даних;
2. поділ ознак на числові та категоріальні;
3. заповнення пропусків у числових ознаках;
4. масштабування числових ознак;
5. заповнення пропусків у категоріальних ознаках;
6. кодування категоріальних ознак;
7. передача підготовлених даних у модель.



На рисунку представлено послідовність попередньої обробки даних перед навчанням моделей машинного навчання.

Важливо, що попередня обробка не виконується окремо від моделі. У системі вона об'єднується з моделлю в один Pipeline. Це зменшує ризик помилок і робить процес навчання більш стабільним.

Тобто модель отримує вже не "сирі" дані, а підготовлені дані у правильному форматі.

Такий підхід особливо корисний для задач, де присутня неповна та невизначена інформація. Саме тому він добре підходить для теми даної дипломної роботи.

2.6 Проєктування алгоритму роботи системи прогнозування

Після визначення структури даних та способу їх попередньої обробки необхідно спроектувати загальний алгоритм роботи системи. Алгоритм має описувати послідовність дій від запуску програми до отримання кінцевого прогнозу.

У розробленій системі передбачено два основні режими роботи:

- консольний запуск програмного модуля;
- робота через графічний інтерфейс користувача.

Консольний режим реалізовано у файлі `main.py`. Він призначений для повного циклу навчання моделей, збереження результатів та формування графіків. Графічний режим реалізовано у файлі `ui.py` і дозволяє користувачу взаємодіяти з системою через віконний інтерфейс.

Загальний алгоритм роботи програмного модуля складається з таких етапів:

1. створення необхідних директорій;
2. генерація синтетичного набору даних;
3. моделювання неповної та невизначеної інформації;
4. поділ даних на навчальну та тестову вибірки;
5. навчання моделей машинного навчання;
6. оцінювання результатів;
7. визначення найкращої моделі;
8. збереження результатів;
9. побудова графіків.

Після запуску програми спочатку створюються службові директорії:

Data

models

plots

У папку `data` зберігається згенерований набір даних та таблиця результатів моделей. У папку `models` записується найкраща навчена модель. У папку `plots` зберігаються графіки, які використовуються для аналізу результатів.

Далі система генерує синтетичний датасет. Кількість записів у поточній реалізації становить 1200. Після цього дані автоматично пошкоджуються: у частині значень створюються пропуски, додається шум і формується категорія unknown.

Це зроблено спеціально, щоб перевірити, як моделі працюють не з ідеальними даними, а з більш наближеними до реальних умов.

Після формування набору даних він поділяється на дві частини:

- навчальну вибірку;
- тестову вибірку.

Навчальна вибірка використовується для побудови моделей, а тестова — для перевірки їх якості. У програмі для тестування використовується 25 % даних.

Після поділу даних система навчає чотири моделі:

- Decision Tree;
- Random Forest;
- Logistic Regression;
- Gaussian Naive Bayes.

Кожна модель проходить однаковий шлях обробки: спочатку дані очищуються та перетворюються, після чого передаються в алгоритм машинного навчання.

Оцінювання моделей виконується за такими показниками:

- accuracy;
- precision;
- recall;
- fl-score;
- classification report;
- confusion matrix.

На основі отриманих результатів система автоматично визначає найкращу модель за показником ассурасу.

Простіше кажучи, програма запускає декілька моделей, порівнює їх між собою і сама вибирає ту, яка дала найкращий результат.

Запуск програми/ Створення директорії/Генерація датасету/Додавання пропусків і шуму/Поділ train/test/Попередня обробка/Навчання моделей/Оцінювання результатів/Вибір найкращої моделі/Збереження результатів

Окремо було спроектовано алгоритм роботи користувача з графічним інтерфейсом.

У цьому режимі користувач виконує такі дії:

1. запускає програму;
2. натискає кнопку «Завантажити дані»;
3. система генерує датасет і навчає моделі;
4. користувач задає вхідні параметри;
5. натискає кнопку «Прогнозувати»;
6. система формує прогноз;
7. результат відображається у правій частині вікна;
8. за потреби результат зберігається у CSV-файл.

2.7 Проєктування інтерфейсу користувача

Під час створення програмного забезпечення важливим етапом є проєктування інтерфейсу користувача. Навіть якщо система містить ефективні алгоритми прогнозування, користувач повинен мати можливість швидко взаємодіяти з програмою та отримувати результати у зрозумілому вигляді.

У межах даної роботи було прийнято рішення реалізувати графічний інтерфейс користувача у вигляді віконного застосунку. Для створення інтерфейсу використовується бібліотека PyQt5, яка дозволяє створювати сучасні графічні системи з підтримкою інтерактивних елементів керування.

Під час проєктування інтерфейсу основна увага приділялась таким вимогам:

- простота використання;
- швидкий доступ до функцій системи;
- логічне розміщення елементів;
- відображення результатів прогнозування;
- можливість перегляду журналу подій;
- підтримка візуалізації результатів.

Інтерфейс системи умовно поділяється на декілька функціональних областей:

- інформаційна панель;
- панель введення параметрів;
- область результатів;
- область графічного відображення;
- журнал подій.

У верхній частині інтерфейсу розташовується інформаційна панель. Вона містить назву системи та службову інформацію щодо поточного стану програми.

Панель використовується для відображення таких повідомлень:

- готовність системи;
- виконання навчання;

- прогнозування;
- повідомлення про помилки.

Основною частиною інтерфейсу є блок введення параметрів. Користувач має можливість задавати характеристики поточної ситуації, які використовуються моделлю для формування прогнозу.

Передбачено введення таких параметрів:

- відстань до супротивника;
- кількість супротивників;
- тип місцевості;
- видимість;
- рівень боєзапасу;
- попередня дія;
- час доби;
- рівень ризику;
- наявність втрат;
- якість зв'язку.

Для введення інформації використовуються різні типи елементів керування:

- списки вибору;
- поля введення;
- повзунки;
- прапорці.

Праворуч від області введення проєктується блок результатів прогнозування.

У цьому блоці повинні відображатися:

- прогнозована дія;
- рівень ймовірності;
- використана модель;

- час виконання прогнозу.

Нижче передбачено область графічної візуалізації, яка використовується для подальшого відображення статистичних матеріалів.

Тут планується показ:

- графіків точності;
- ймовірностей класів;
- результатів аналізу моделей.

У нижній частині інтерфейсу передбачається реалізація журналу подій.

Журнал використовується для інформування користувача щодо роботи системи та відображає:

- запуск процесів;
- виконання навчання;
- результати прогнозування;
- помилки.

Назва системи	Статус
Вхідні параметри	Результати прогнозування
	Графічна візуалізація
Журнал подій	

На схемі представлено загальну структуру графічного інтерфейсу системи прогнозування.

При проектуванні було використано принцип розділення функціональних областей. Це дозволяє зробити інтерфейс більш зрозумілим та зменшує кількість дій, необхідних для роботи користувача із системою.

Таким чином, спроектований інтерфейс забезпечує зручну взаємодію користувача з програмною системою та створює основу для подальшої реалізації функціональних можливостей.

Висновки до розділу 2

У другому розділі було виконано вибір засобів реалізації програмної системи та спроектовано її архітектуру. Було обґрунтовано використання мови програмування Python та бібліотек машинного навчання, аналізу даних і створення графічного інтерфейсу.

Також було розроблено структуру синтетичного набору даних, визначено механізм попередньої обробки інформації та спроектовано алгоритм роботи системи.

Окрему увагу приділено розробці інтерфейсу користувача, що забезпечує зручну взаємодію із системою та відображення результатів прогнозування

РОЗДІЛ 3

РЕАЛІЗАЦІЯ ПРОГРАМНОГО МОДУЛЯ ТА ЕКСПЕРИМЕНТАЛЬНЕ ДОСЛІДЖЕННЯ

3.1 Реалізація генерації синтетичного набору даних

Після завершення етапу проєктування було реалізовано програмний модуль формування синтетичного набору даних. Основною метою цього етапу є створення навчальної вибірки, яка використовується для побудови моделей машинного навчання.

У розробленій системі генерація даних реалізується автоматично. Для цього використовується окрема функція створення набору даних.

Кожний запис у наборі даних описує окрему умовну ситуацію та містить характеристики середовища, параметри супротивника і цільову змінну.

Для формування набору використовується механізм випадкової генерації параметрів:

- відстань до супротивника;
- кількість супротивників;
- місцевість;
- рівень видимості;
- рівень ризику;
- втрати;
- попередні дії;
- час доби;
- якість зв'язку.

Після створення основних параметрів система виконує розрахунок імовірностей майбутньої дії супротивника.

Для цього використовується набір умовних правил. Наприклад:

- невелика відстань та високий боєзапас підвищують ймовірність атаки;

- великі втрати збільшують шанс відступу;
- низька видимість підвищує ймовірність розвідки.

Це дозволяє створити не випадковий набір значень, а сформувати певні статистичні закономірності.

У програмному коді генерація набору даних виконується таким фрагментом:

```
dataset = generate_synthetic_dataset(rows_count=1200)
```

У поточній реалізації формується 1200 записів.

Якщо пояснювати простіше, програма створює великий набір умовних ситуацій, які потім використовуються як "досвід" для навчання моделей.

Після завершення генерації дані автоматично зберігаються:

data/dataset.csv

```
distance_to_enemy,enemy_count,terrain_type,visibility_level,ammo_level,previous_enemy_action,time_of_day,risk_level,losses_detected,communication_quality,next_enemy_action
96.234122131382,9.0,urban,90.06722717453994,52.5295663659051,retreat,morning,100.0,61.52221358174107,77.74639853385835,retreat
47.10235172912788,7.0,forest,84.1590898494201,61.397639648930536,reconnaissance,night,63.868912103250445,10.436871761915981,54.08129213620916,attack
96.9375755439709,0.0,mountain,82.23168674241305,39.93653845377369,attack,night,,50.48784945484598,,reconnaissance
4.932677324139143,7.0,plain,73.6493982840213,71.57726331427308,retreat,,38.43364068767575,22.793167247087496,51.02995554536547,attack
12.053151646225256,,urban,28.99868171502832,76.87391237549284,waiting,day,80.87755389407103,,36.314465651637875,waiting
97.83143494978283,,forest,39.09614937328656,,flanking,morning,19.357620004069243,0.43667544996525076,82.82566241272573,flanking
65.54978536358001,4.0,mountain,48.747836552319754,67.7316881093052,attack,morning,15.103427724757816,51.58106246634255,49.820943478933756,reconnaissance
90.45096074749337,6.0,mountain,55.13123086992862,65.13594358865579,flanking,day,0.0,32.87728781835268,,reconnaissance
102.4190336230393,0.0,forest,10.053329835536687,36.54868412359539,waiting,day,78.06094201379106,39.29645501640663,61.18016487774289,flanking
40.86027669436001,9.0,desert,29.13666082777396,0.0,attack,morning,53.91438761068818,33.71637349299885,22.350318934865562,reconnaissance
34.62164158360944,9.0,mountain,46.603173129999526,38.23924795343386,retreat,day,60.366729326001945,28.446837540553226,12.671824692461913,attack
110.86256442481566,4.0,desert,68.82227599726606,36.36918404581047,flanking,night,85.08473382518535,41.77378350998172,37.05806000100894,retreat
14.185086707933724,6.0,desert,49.7390516987018,22.41509427758734,flanking,day,50.31552006867121,,87.90608663641068,flanking
83.00048662304012,10.0,urban,62.59375609352176,62.8446761312653,attack,evening,6.283092093388412,24.97445343751066,9.344925234831035,reconnaissance
```

На рисунку представлено фрагмент автоматично сформованого набору даних, який використовується для навчання моделей машинного навчання.

Після аналізу отриманих результатів встановлено, що система успішно формує набір даних із необхідними параметрами та забезпечує створення навчального середовища для подальших експериментів.

3.2 Реалізація механізму пошкодження даних та моделювання невизначеності

Однією з головних особливостей розробленої системи є можливість роботи в умовах неповної та невизначеної інформації. Для цього у програмному модулі реалізовано спеціальний механізм штучного пошкодження даних.

Його основне завдання полягає у створенні умов, максимально наближених до реальних ситуацій, коли інформація може бути неповною, неточною або містити випадкові спотворення.

У реальних умовах інформаційні системи часто стикаються з такими проблемами:

- відсутність частини параметрів;
- помилки передачі інформації;
- випадкові похибки;
- затримка отримання даних;
- суперечливість інформації.

Через це навчання моделей лише на ідеальних наборах даних може призводити до погіршення результатів у реальному використанні.

У програмі для моделювання невизначеності реалізовано окрему функцію:

```
damage_dataset(dataset)
```

Функція виконує декілька послідовних етапів пошкодження даних.

На першому етапі створюються пропущені значення.

У кожній ознаці випадковим чином обирається частина записів, після чого їм присвоюється значення:

```
np.nan
```

У поточній реалізації приблизно 8 % записів містять пропущені значення.
Якщо говорити простіше, система спеціально "псує" частину інформації.
Наприклад:

Було:

$$\begin{aligned} enemy_count &= 10 \\ visibility_level &= 75 \end{aligned}$$

Після пошкодження:

$$\begin{aligned} enemy_count &= NaN \\ visibility_level &= NaN \end{aligned}$$

Другим етапом є додавання випадкового шуму.

Для числових параметрів створюється випадкове відхилення, яке додається до початкового значення.

Наприклад:

Було:

$$distance_to_enemy = 40$$

Після додавання шуму:

$$distance_to_enemy = 47$$

або:

$$distance_to_enemy = 35$$

Це дозволяє моделювати ситуації, коли система отримує не зовсім точну інформацію.

Додавання шуму реалізується за допомогою нормального розподілу випадкових величин.

Крім цього, у системі створюються невизначені категорії.

Для цього частині записів присвоюється значення:

$$unknown$$

Наприклад:

previous_enemy_action = unknown

Такий підхід дозволяє імітувати ситуації, коли інформація про попередню поведінку відсутня.

Після завершення пошкодження система додатково перевіряє діапазони значень та виконує обмеження параметрів.

Наприклад:

- відстань не може бути меншою за 0;
- ризик не може перевищувати 100;
- кількість супротивників не може бути від’ємною.

Це дозволяє уникнути появи некоректних значень.

```
,11.0,mountain,8.992113341443144,83.84022314056666,waiting,evening,10.195734522110698,,67.92989008910538,reconnaissance
,15.0,,61.78073085929728,32.009155946806786,,evening,94.70669703809558,45.80461323024748,67.1909421235399,retreat
90.79088833730027,4.0,desert,22.066648735736617,8.690352835185797,flanking,evening,34.6736749892157,45.072606162405535,20.853268528464717,retreat
,14.0,urban,64.99346286681234,72.2680883542599,unknown,evening,81.76302521071808,38.48373506696673,56.87277796300501,reconnaissance
22.954418574946693,2.0,forest,49.14755772640851,78.0513131863888,reconnaissance,evening,11.43763055444477,,16.58798884830302,flanking
96.81521025361766,6.0,urban,66.85692694595673,24.41525773767954,flanking,day,78.41180970170137,,57.8663221997018,retreat
7.419129054648988,11.0,plain,31.539111638635468,36.464051045098806,waiting,night,38.333915342617686,10.540407354331514,16.73518967331917,reconnaissance
101.07029341021074,11.0,forest,72.0659954008218,,reconnaissance,morning,58.59448391332467,66.71763094183825,37.03771939198887,attack
110.14854317111605,11.0,desert,50.58195473658384,58.748621536319334,flanking,morning,87.82143877546365,36.60610262823593,0.0,flanking
121.0068855305577,7.0,desert,35.90327657698866,,reconnaissance,night,97.45616507409805,38.08524947348765,68.52002269696608,reconnaissance
45.805636913951915,15.0,plain,11.952959245884973,22.233877227786685,retreat,night,6.6556970930612955,,55.369527346340774,retreat
,5.0,urban,44.37902826857292,9.93080646284278,retreat,night,16.331638908704832,64.59626845949757,78.11659600839144,waiting
75.11957997981094,,mountain,77.80546509459758,75.3786802297352,unknown,morning,100.0,44.17358775974897,80.18275193216706,retreat
6.907568983796638,12.0,urban,55.2283519698354,39.19351153257876,reconnaissance,day,14.428617397570871,44.826112059140414,32.43187078517976,waiting
14.968592840416952,9.0,forest,27.9611278111367,61.00639085917726,retreat,day,,63.15616793945881,16.39729221448689,flanking
0.0,3.0,forest,54.4795000094559,12.987054359585922,waiting,evening,79.88782966170868,49.54395425601881,48.84051973719482,flanking
52.00588615572992,,plain,17.752083913068866,,attack,night,0.0,37.98651094848374,85.3642251619107,waiting
,5.0,mountain,71.65414581227061,80.3410610042774,waiting,morning,100.0,3.277297851943472,93.28446709648199,attack
15.762150598510702,8.0,plain,53.40036360839518,29.070358620422393,retreat,morning,59.6817031512921,28.651050660167144,,waiting
77.76796429172691,16.0,,41.42243334333484,79.59612867720068,unknown,morning,38.685734212544666,68.99686328617913,97.22034726768555,retreat
76.60628474794496,6.0,urban,78.05996370758388,0.0,attack,morning,58.3328027832134,10.65272565993874,,retreat
116.16273201562232,,desert,53.93063978922541,87.70110660724181,attack,day,58.72382745236229,15.808012076912217,86.87777350392341,flanking
,7.0,desert,100.0,14.469018824509813,reconnaissance,night,55.4461457599341,25.86357897314238,47.19698298707155,retreat
40.54166669695192,10.0,,25.8223970434925,60.25805858395033,,night,41.83705662118264,35.75162478262625,35.61527663427218,reconnaissance
27.374897001999532,5.0,,33.689289521793704,68.19472000910667,waiting,night,56.59775285056721,68.94491438311903,24.24758784060337,waiting
67.6409048975544,14.0,mountain,70.73061547861965,100.0,attack,,46.249921166214925,49.982883815839884,85.668848564894,retreat
102.50484304080925,6.0,forest,53.85172981359656,5.139731176572734,flanking,morning,90.73037173561828,26.36135325582246,48.021934790489375,retreat
60.52951276645746,11.0,desert,98.45647530117452,5.85958760584111,flanking,,38.80167430646615,36.626500671457364,36.088548184212314,retreat
62.184115027732915,7.0,urban,,39.93298058052542,flanking,night,72.1365890485323,24.661834465442357,80.48484067486163,waiting
8.33421143892694,1.0,urban,31.637896327167304,68.41562600398565,retreat,morning,75.93251779467387,23.09989050098446,47.8935192009725,waiting
48.42155325960628,1.0,desert,97.12382913323664,98.19051821538345,unknown,morning,81.42583436079627,3.098733251260636,,flanking
75.22839176967854,11.0,urban,38.14448429376881,85.02598812814045,reconnaissance,evening,33.081890015919114,61.890324337997875,18.70645923024542,attack
80.7993942190791,11.0,desert,39.65439489851036,33.197580126802244,waiting,morning,97.19509761394659,22.489839474053646,61.563385654939005,waiting
6.240997320436409,9.0,urban,60.88727342529762,70.22002323309522,retreat,day,,13.40817081578118,73.93027086858959,waiting
26.664294709002068,4.0,forest,91.63287681333993,5.509354992842036,flanking,day,65.3251974998872,27.484400457755783,28.933773354577127,retreat
34.40974087411493,0.0,urban,73.30313390281343,49.75794872908645,waiting,morning,71.88175115375839,15.824735558349847,78.27526571017944,retreat
98.68947422064399,5.0,forest,22.034907191839782,39.0040365174854,waiting,morning,21.440853149931574,36.299655632120576,75.51618643085439,attack
113.3220852971466,6.0,mountain,73.36107810731724,71.75298456871244,unknown,evening,57.09694754758149,45.311761686287305,92.77939603938474,attack
38.51017514490978,0.0,,85.141977198343,46.354853267230105,attack,evening,83.32770269692325,63.97964502094744,24.67670035352131,waiting
129.63857257586184,5.0,,72.66963304006659,35.89929841004823,flanking,evening,40.757763782997536,14.035373175280277,25.30366408424425,attack
7.131575864517496,0.0,mountain,35.37621195186146,27.40381112160857,flanking,evening,47.20372681505505,45.1147110016163,84.31326194482955,retreat
96.90219547269412,11.0,plain,78.61032283577786,100.0,unknown,night,71.67585809653617,41.999596776288286,23.180145044584854,reconnaissance
100.80633772782775,6.0,desert,60.00211872385222,17.888241290893834,reconnaissance,evening,14.134758815328784,26.596394027957743,33.3750677709059,flanking
32.1597046849155,15.0,plain,13.147707292587224,77.80124964894141,retreat,evening,14.495856036882891,0.0,attack
77.5226424528106,5.0,urban,100.0,40.79703145973859,flanking,evening,13.763453056103963,61.49188349183794,100.0,reconnaissance
```

На рисунку представлено фрагмент набору даних після застосування механізму моделювання неповної та невизначеної інформації.

Проведене дослідження показало, що реалізований механізм пошкодження даних дозволяє формувати складні умови для навчання моделей машинного навчання та забезпечує більш реалістичне середовище тестування.

3.3 Реалізація попередньої обробки даних

Після формування синтетичного набору даних та моделювання невизначеності наступним важливим етапом є попередня обробка інформації. Як було зазначено раніше, моделі машинного навчання погано працюють із пропущеними значеннями, текстовими параметрами та даними, що мають різні масштаби.

Тому перед початком навчання у системі реалізовано автоматичний механізм підготовки інформації.

У розробленій системі попередня обробка виконується через конвеєр (Pipeline), який поєднує всі етапи роботи з даними в єдиний процес.

Основна перевага такого підходу полягає в тому, що система самостійно виконує необхідні перетворення та передає вже підготовлені дані у модель.

У програмному коді створення конвеєра виглядає таким чином:

```
Pipeline(  
    steps=[  
        ("preprocessor", build_preprocessor()),  
        ("model", model),  
    ]  
)
```

Попередня обробка складається з декількох етапів.

На першому етапі система виконує пошук пропущених значень.

Для числових параметрів використовується метод:

```
SimpleImputer(strategy="median")
```

У цьому випадку пропущені значення автоматично замінюються медіаною.

Наприклад:

До обробки:

distance_to_enemy = NaN

Після обробки:

distance_to_enemy = 54.3

Для категоріальних ознак використовується інший механізм:

SimpleImputer(strategy="most_frequent")

У цьому випадку пропущені значення замінюються найбільш поширеною категорією.

Наприклад:

Було:

terrain_type = NaN

Стало:

terrain_type = forest

Після заповнення пропусків система виконує масштабування числових параметрів.

Для цього використовується:

StandardScaler()

Це необхідно для приведення всіх параметрів до приблизно однакового масштабу.

Якщо пояснити простіше, система вирівнює дані, щоб великі числа не мали надто сильного впливу на результати навчання.

Наприклад:

Було:

distance = 120

risk = 10

Після масштабування:

$$\begin{aligned} distance &= 1.2 \\ risk &= -0.4 \end{aligned}$$

Для роботи з текстовими параметрами використовується:

OneHotEncoder()

Оскільки моделі машинного навчання не можуть напряму працювати з текстом, категорії переводяться у числовий формат.

Наприклад:

Було:

$$terrain = forest$$

Після перетворення:

$$\begin{aligned} terrain_forest &= 1 \\ terrain_urban &= 0 \\ terrain_plain &= 0 \end{aligned}$$

Усі ці операції автоматично об'єднуються через:

ColumnTransformer()

Після цього система формує фінальний набір ознак та передає його моделі машинного навчання.

```

model,accuracy,precision,recall,f1_score,classification_report,confusion_matrix
Random Forest,0.32,0.31037189244433144,0.32,0.303761150585458,"
precision    recall  f1-score   support

   attack      0.29      0.22      0.25        49
  flanking      0.27      0.14      0.19        42
reconnaissance  0.34      0.38      0.36        61
   retreat      0.34      0.51      0.40        91
   waiting      0.29      0.18      0.22        57

 accuracy      0.32        300
macro avg      0.30      0.29      0.28        300
weighted avg   0.31      0.32      0.30        300
", "[[11, 15, 11, 3, 9], [9, 46, 18, 9, 9], [7, 26, 23, 2, 3], [4, 19, 9, 6, 4], [7, 31, 7, 2, 10]]"
Logistic Regression,0.31,0.3274422017010528,0.31,0.3066097208516072,"
precision    recall  f1-score   support

   attack      0.25      0.35      0.29        49
  flanking      0.19      0.29      0.23        42
reconnaissance  0.35      0.41      0.38        61
   retreat      0.36      0.18      0.24        91
   waiting      0.42      0.40      0.41        57

 accuracy      0.31        300
macro avg      0.31      0.32      0.31        300
weighted avg   0.33      0.31      0.31        300
", "[[17, 8, 8, 12, 4], [21, 16, 21, 19, 14], [13, 8, 25, 9, 6], [9, 3, 10, 12, 8], [7, 10, 7, 10, 23]]"
Decision Tree,0.30333333333333334,0.2800326694499205,0.30333333333333334,0.2829413407533843,"
precision    recall  f1-score   support

   attack      0.19      0.12      0.15        49
  flanking      0.17      0.10      0.12        42
reconnaissance  0.29      0.25      0.27        61
   retreat      0.36      0.54      0.43        91
   waiting      0.31      0.30      0.30        57

 accuracy      0.30        300
macro avg      0.26      0.26      0.25        300
weighted avg   0.28      0.30      0.28        300
", "[[6, 19, 7, 7, 10], [7, 49, 17, 7, 11], [8, 25, 15, 5, 8], [4, 18, 7, 4, 9], [6, 27, 6, 1, 17]]"
Gaussian Naive Bayes,0.2833333333333333,0.29526575668236976,0.2833333333333333,0.28743342723280063,"
precision    recall  f1-score   support

   attack      0.22      0.24      0.23        49
  flanking      0.15      0.21      0.18        42
reconnaissance  0.35      0.33      0.34        61
   retreat      0.35      0.33      0.34        91
   waiting      0.32      0.25      0.28        57

 accuracy      0.28        300
macro avg      0.27      0.27      0.27        300
weighted avg   0.30      0.28      0.29        300
", "[[12, 15, 9, 9, 4], [14, 30, 12, 20, 15], [14, 9, 20, 10, 8], [7, 14, 9, 9, 3], [7, 18, 7, 11, 14]]"

```

На рисунку представлено результат успішного виконання етапу попередньої обробки даних перед навчанням моделей машинного навчання.

Таким чином, реалізований механізм попередньої обробки забезпечує автоматичне очищення, перетворення та підготовку інформації, що дозволяє підвищити якість навчання моделей.

3.4 Реалізація моделей машинного навчання

Після завершення попередньої обробки даних система переходить до етапу навчання моделей машинного навчання. Основною метою цього етапу є побудова декількох моделей та подальше визначення найбільш ефективного алгоритму прогнозування.

У межах роботи реалізовано чотири моделі:

- дерево рішень (Decision Tree);
- випадковий ліс (Random Forest);
- логістична регресія (Logistic Regression);
- наївний байєсівський класифікатор (Gaussian Naive Bayes).

Набір моделей формується окремою функцією:

```
def get_models():  
    return {  
        "Decision Tree": ...,  
        "Random Forest": ...,  
        "Logistic Regression": ...,  
        "Gaussian Naive Bayes": ...  
    }
```

Усі моделі використовують однаковий механізм підготовки даних та працюють через спільний конвеєр обробки.

Такий підхід дозволяє проводити коректне порівняння результатів.

Реалізація дерева рішень

Першою моделлю є дерево рішень.

Для реалізації використовується:

```
DecisionTreeClassifier(  
    random_state=42,  
    max_depth=8  
)
```

Для обмеження перенавчання використовується параметр:

Max_depth=8

Це означає, що модель не може створювати надто глибоку структуру дерева.

Якщо пояснити простими словами, це зроблено для того, щоб модель не "запам'ятовувала" навчальний набір занадто точно.

Реалізація Random Forest

Другою моделлю є випадковий ліс.

У програмі використовується:

RandomForestClassifier(

n_estimators=160,
random_state=42,
max_depth=10,
class_weight="balanced"

)

У системі використовується:

- 160 дерев;
- глибина 10;
- балансування класів.

Використання ансамблю дерев дозволяє зменшити вплив шуму та покращити точність прогнозування.

Саме ця модель очікувано має показати найкращі результати.

Реалізація логістичної регресії

Наступною моделлю є логістична регресія:

LogisticRegression(

random_state=42,
max_iter=1000,
class_weight="balanced"

)

Збільшення кількості ітерацій використовується для покращення процесу навчання.

Модель має простішу структуру та працює швидше за ансамблеві алгоритми.

Реалізація найвнього байєсівського класифікатора

Четвертою моделлю є:

GaussianNB()

Основною перевагою цієї моделі є швидкість роботи.

Проте модель базується на припущенні незалежності ознак, що не завжди відповідає реальним умовам.

Через це точність прогнозування може бути нижчою.

Після створення всіх моделей система автоматично запускає процес навчання.

Для кожної моделі виконуються:

- навчання;
- прогноз;
- обчислення метрик;
- побудова матриці помилок;
- збереження результатів.

```

=== Порівняння моделей ===
      model  accuracy  precision  recall  f1_score
Random Forest  0.320000  0.310372  0.320000  0.303761
Logistic Regression  0.310000  0.327442  0.310000  0.306610
Decision Tree  0.303333  0.280033  0.303333  0.282941
Gaussian Naive Bayes  0.283333  0.295266  0.283333  0.287433

```

```

=== Детальні звіти класифікації ===

```

```

--- Random Forest ---
      precision  recall  f1-score  support

  attack      0.29    0.22    0.25     49
  flanking     0.27    0.14    0.19     42
reconnaissance 0.34    0.38    0.36     61
  retreat     0.34    0.51    0.40     91
  waiting     0.29    0.18    0.22     57

  accuracy                0.32     300
  macro avg      0.30    0.29    0.28     300
  weighted avg   0.31    0.32    0.30     300

```

```

Confusion matrix:
[[11 15 11  3  9]
 [ 9 46 18  9  9]
 [ 7 26 23  2  3]
 [ 4 19  9  6  4]
 [ 7 31  7  2 10]]

```

```

--- Logistic Regression ---
      precision  recall  f1-score  support

  attack      0.25    0.35    0.29     49
  flanking     0.19    0.29    0.23     42
reconnaissance 0.35    0.41    0.38     61
  retreat     0.36    0.18    0.24     91
  waiting     0.42    0.40    0.41     57

  accuracy                0.31     300
  macro avg      0.31    0.32    0.31     300
  weighted avg   0.33    0.31    0.31     300

```

```

Confusion matrix:
[[17  8  8 12  4]
 [21 16 21 19 14]
 [13  8 25  9  6]
 [ 9  3 10 12  8]
 [ 7 10  7 10 23]]

```

```

--- Decision Tree ---
              precision    recall  f1-score   support

   attack         0.19      0.12      0.15         49
  flanking         0.17      0.10      0.12         42
reconnaissance     0.29      0.25      0.27         61
   retreat         0.36      0.54      0.43         91
   waiting         0.31      0.30      0.30         57

 accuracy         0.30         300
 macro avg         0.26      0.26      0.25         300
 weighted avg         0.28      0.30      0.28         300

Confusion matrix:
[[ 6 19  7  7 10]
 [ 7 49 17  7 11]
 [ 8 25 15  5  8]
 [ 4 18  7  4  9]
 [ 6 27  6  1 17]]

--- Gaussian Naive Bayes ---
              precision    recall  f1-score   support

   attack         0.22      0.24      0.23         49
  flanking         0.15      0.21      0.18         42
reconnaissance     0.35      0.33      0.34         61
   retreat         0.35      0.33      0.34         91
   waiting         0.32      0.25      0.28         57

 accuracy         0.28         300
 macro avg         0.28      0.27      0.27         300
 weighted avg         0.30      0.28      0.29         300

Confusion matrix:
[[12 15  9  9  4]
 [14 30 12 20 15]
 [14  9 20 10  8]
 [ 7 14  9  9  3]
 [ 7 18  7 11 14]]

```

На Рисуноках показано Виведення інформації про навчання моделей

Після завершення навчання система переходить до оцінювання ефективності моделей та визначення найбільш точного алгоритму.

3.5 Навчання та тестування моделей

Після реалізації моделей машинного навчання було виконано їх навчання та подальше тестування. Основною метою цього етапу є оцінювання ефективності алгоритмів та визначення моделі, яка демонструє найкращі результати прогнозування.

Перед початком навчання набір даних автоматично поділяється на дві частини:

- навчальну вибірку;
- тестову вибірку.

У програмному коді використовується:

```
x_train, x_test, y_train, y_test =  
train_test_split(  
    x,  
    y,  
    test_size=0.25,  
    random_state=42,  
    stratify=y  
)
```

Для навчання використовується 75 % даних, а решта 25 % застосовується для тестування моделей.

Поділ набору даних дозволяє оцінити, наскільки добре модель працює не лише на навчальних даних, а й на нових прикладах.

Якщо пояснювати простими словами, система навчається на одній частині інформації, а потім перевіряється на даних, які вона раніше не бачила.

Після завершення навчання система виконує оцінювання моделей.

У роботі використовуються такі показники:

- Accuracy;
- Precision;
- Recall;
- F1-score;

- Classification Report;
- Confusion Matrix.

Показник Accuracy характеризує загальну точність моделі.

Precision визначає частку правильних позитивних результатів.

Recall показує здатність моделі знаходити правильні варіанти серед усіх можливих.

F1-score використовується для комплексного оцінювання якості.

Після розрахунку метрик система автоматично сортує результати та визначає модель із найвищим значенням accuracy.

```
results_df.sort_values(  
  
by="accuracy",  
ascending=False  
)
```

Таким чином система не потребує ручного аналізу результатів.

```
=== Порівняння моделей ===
```

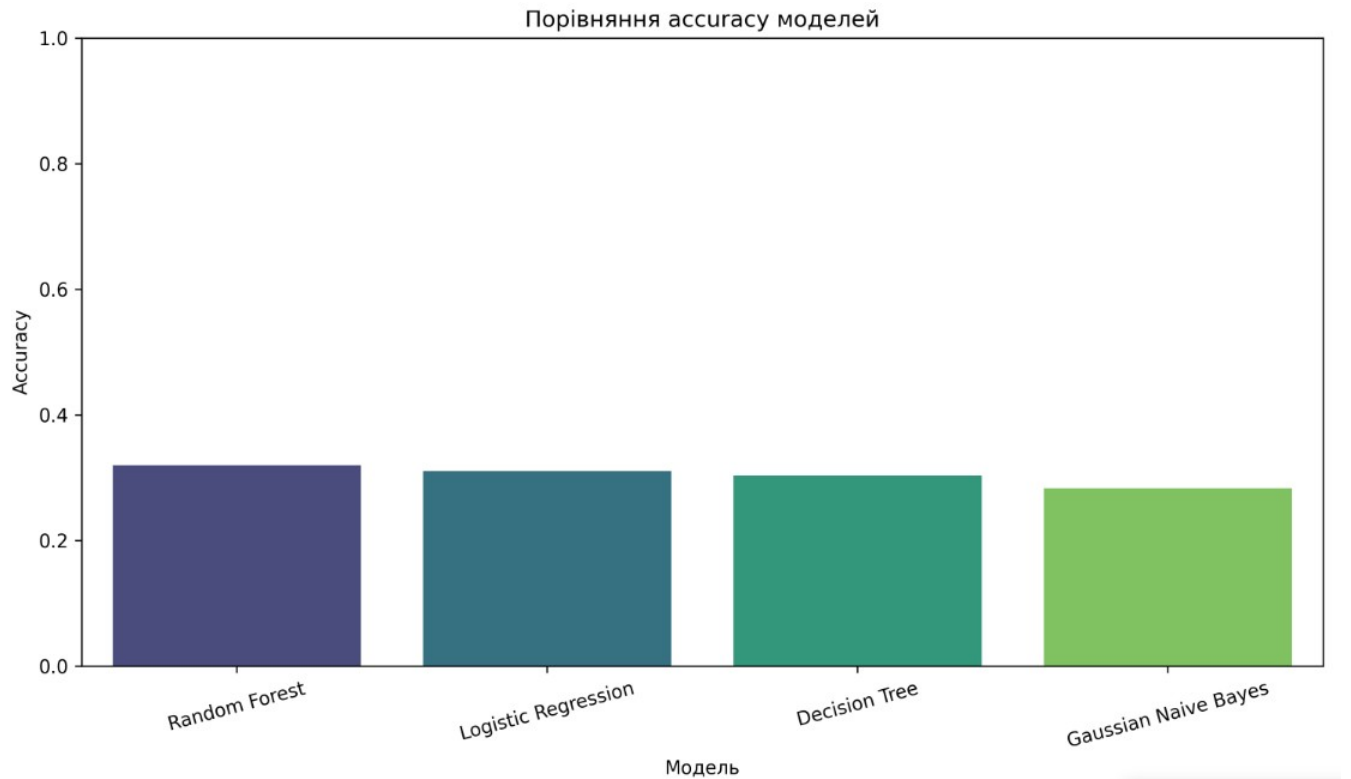
model	accuracy	precision	recall	f1_score
Random Forest	0.320000	0.310372	0.320000	0.303761
Logistic Regression	0.310000	0.327442	0.310000	0.306610
Decision Tree	0.303333	0.280033	0.303333	0.282941
Gaussian Naive Bayes	0.283333	0.295266	0.283333	0.287433

На рисунку представлено результати порівняння моделей машинного навчання після завершення навчання.

Для візуального аналізу результатів система автоматично будує графік порівняння точності моделей.

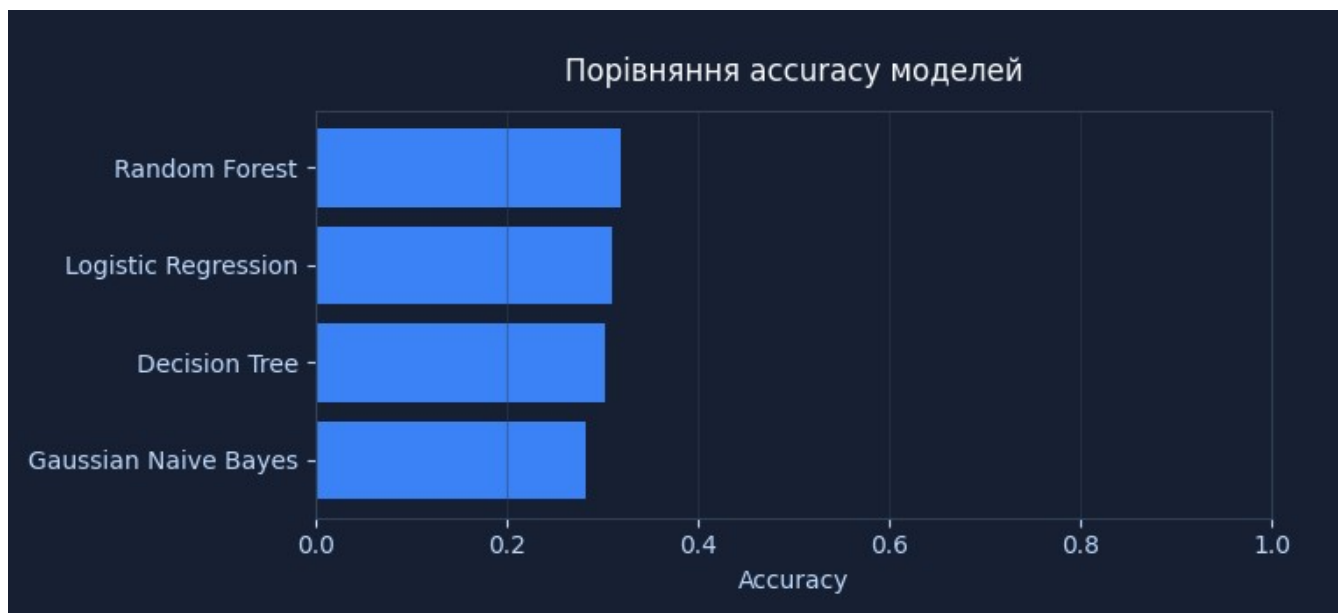
Графік формується після завершення процесу навчання та зберігається у директорію:

plots/



На рисунку представлено графік порівняння точності моделей машинного навчання. Після аналізу результатів необхідно визначити модель із найкращими показниками. У поточній реалізації система автоматично обирає алгоритм із найбільшим значенням ассигасу та зберігає його для подальшого використання.

best_model.pkl



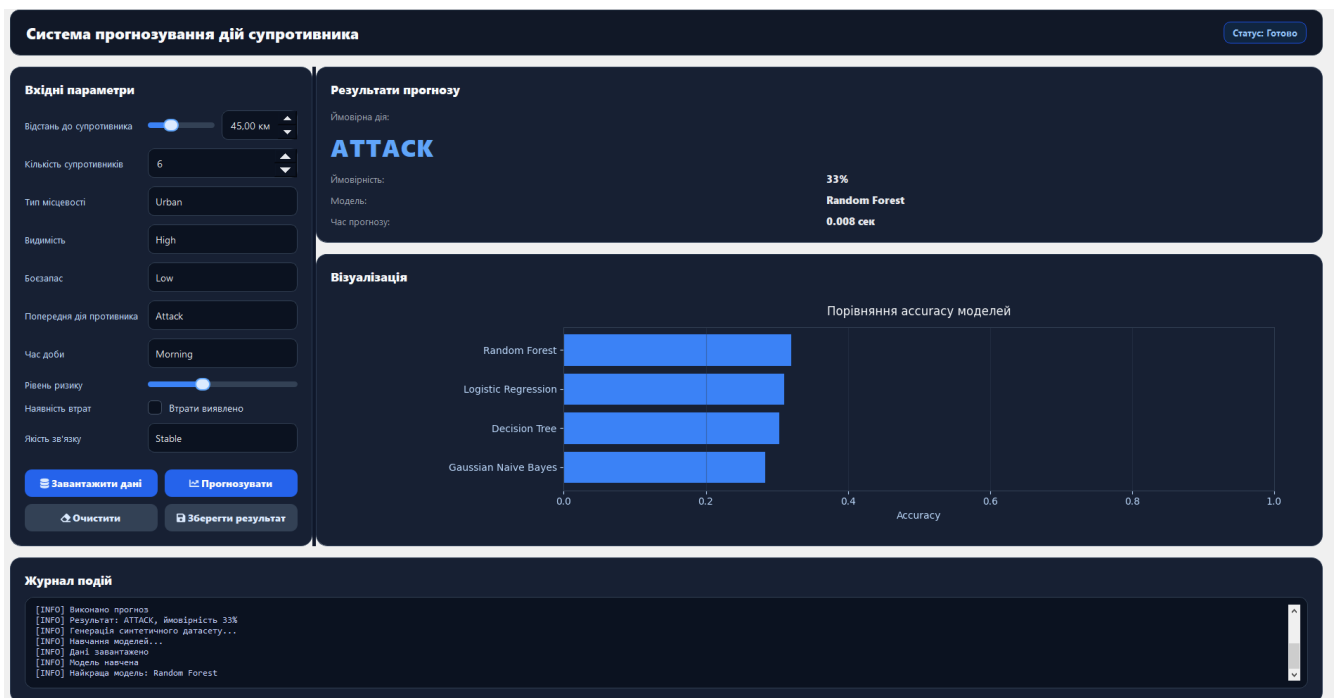
3.6 Реалізація графічного інтерфейсу та аналіз результатів прогнозування

Після реалізації алгоритмів машинного навчання було створено графічний інтерфейс користувача, який забезпечує взаємодію із системою прогнозування. Основною метою інтерфейсу є спрощення роботи користувача та надання можливості отримувати результати без необхідності запуску команд вручну.

Розроблений інтерфейс реалізовано засобами бібліотеки PyQt5. Його структура складається з декількох функціональних блоків:

- інформаційна панель;
- панель введення параметрів;
- область результатів;
- графічна область;
- журнал подій.

Після запуску застосунку користувач отримує доступ до головного вікна системи.



На рисунку представлено головне вікно розробленої системи прогнозування.

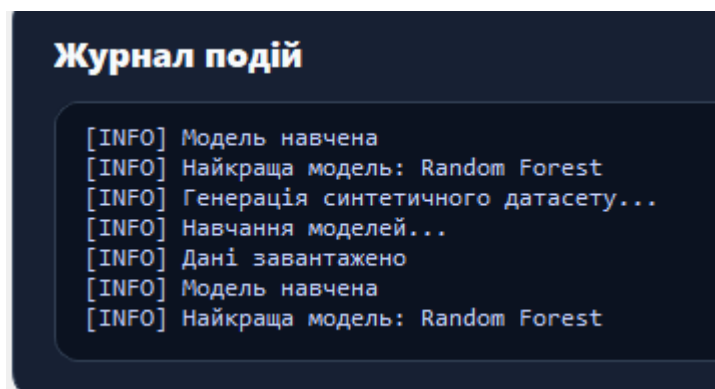
Для запуску процесу навчання користувач натискає кнопку:

Завантажити дані

Після цього система автоматично:

- генерує синтетичний датасет;
- запускає навчання моделей;
- оцінює результати;
- обирає найкращу модель;
- зберігає результати.

Після завершення процесу в журналі відображається інформація про успішне виконання операцій.



На рисунку представлено журнал подій після завершення навчання моделей.

Після завершення навчання користувач може вводити параметри ситуації та запускати прогнозування.

Для перевірки роботи системи рекомендується використати такі параметри:

Відстань: 25

Кількість супротивників: 10

Місцевість: Forest

Видимість: Low

Боезапас: High

Попередня дія: Reconnaissance

Час: Night

Ризик: 35

Втрати: відсутні

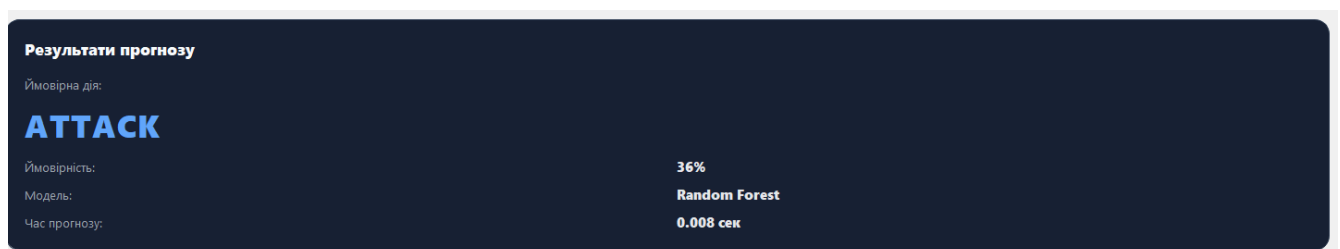
Якість зв'язку: Stable

Після введення параметрів необхідно натиснути:

Прогнозувати

Після виконання прогнозу система відображає:

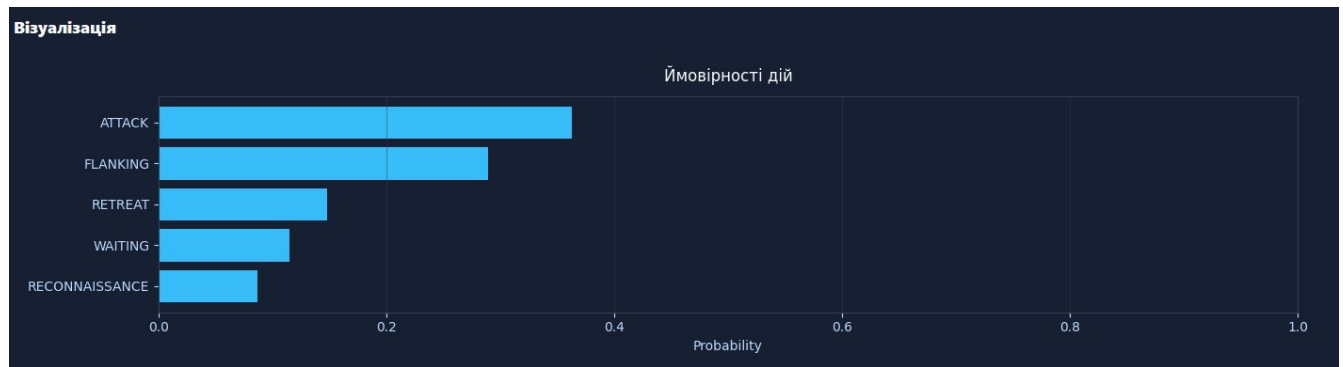
- дію супротивника;
- імовірність;
- модель;
- час виконання.



На рисунку представлено результат прогнозування дії супротивника на основі введених параметрів.

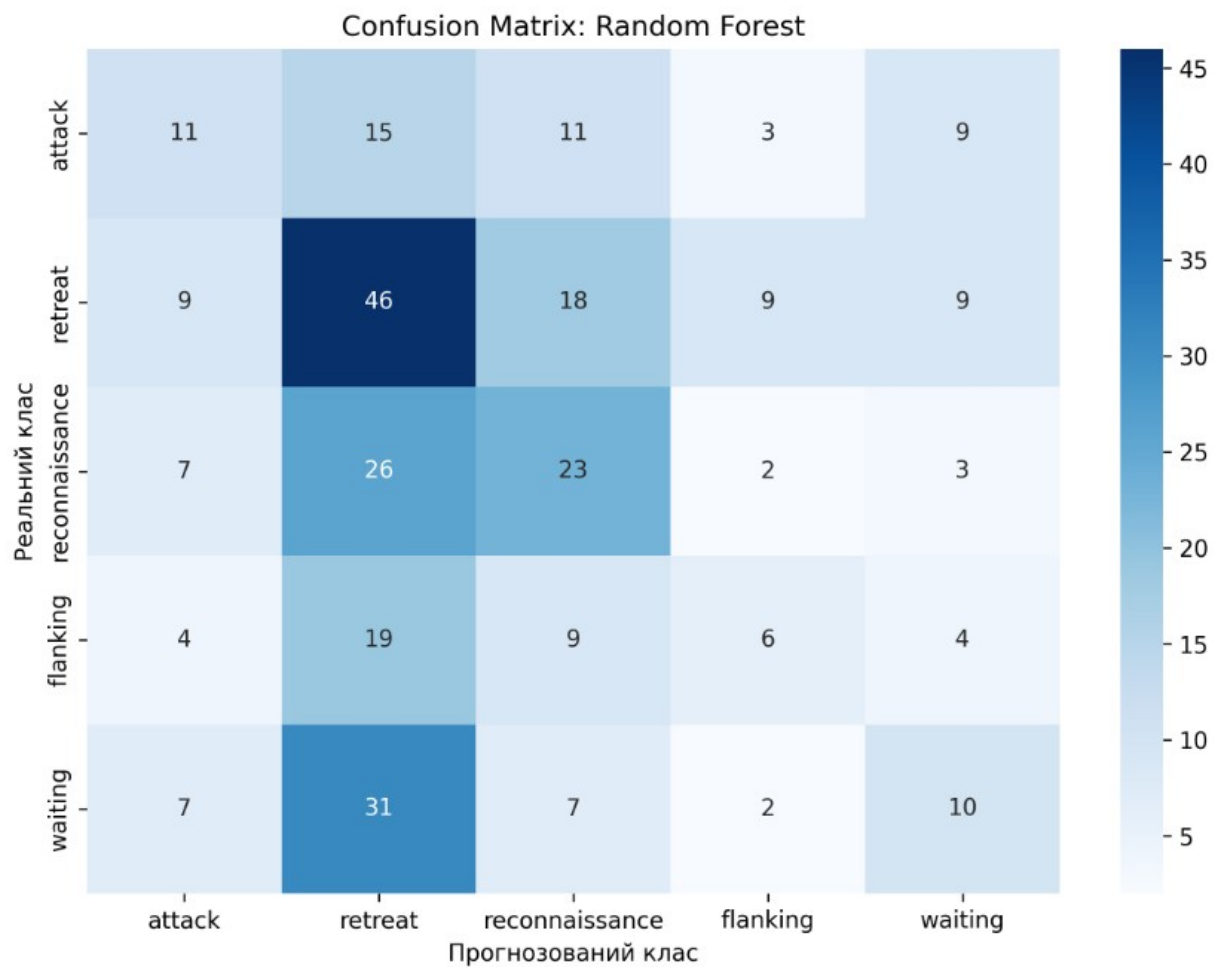
Крім текстового результату система формує графічне представлення ймовірностей для всіх можливих класів.

Це дозволяє оцінити не лише кінцевий результат, але й зрозуміти, наскільки впевнено модель прийняла рішення.

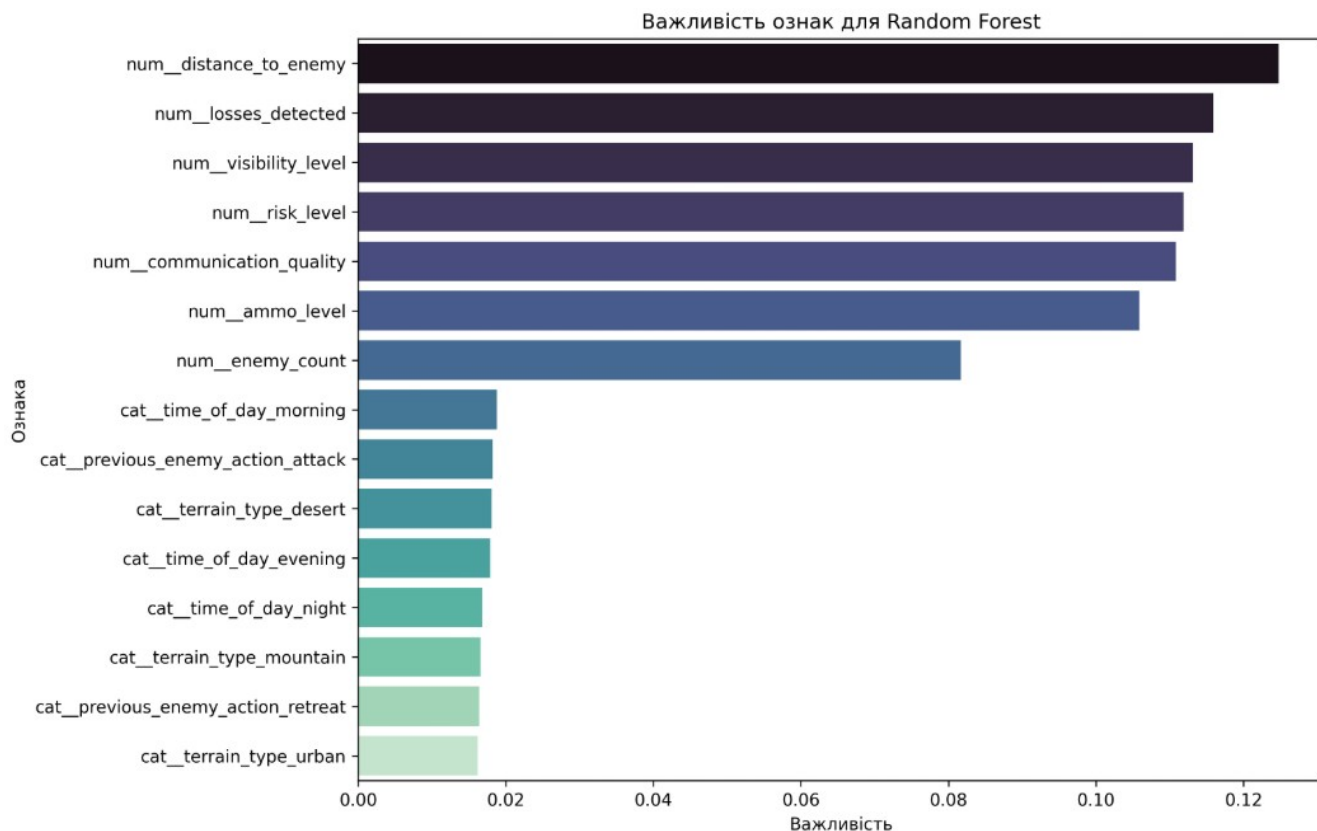


На рисунку представлено графічне відображення ймовірностей для можливих дій супротивника.

Для детального аналізу роботи моделей система також автоматично будує матрицю помилок.



Також система формує графік важливості ознак.



Після проведення експериментів встановлено, що розроблена система успішно виконує генерацію даних, навчання моделей та прогнозування. Інтерфейс забезпечує зручну взаємодію користувача із системою та дозволяє наочно аналізувати результати роботи.

3.7 Аналіз результатів експериментального дослідження

Після реалізації програмного модуля та проведення серії експериментів було виконано аналіз отриманих результатів. Основною метою дослідження було визначення ефективності різних моделей машинного навчання при прогнозуванні дій супротивника в умовах неповної та невизначеної інформації.

Під час експерименту всі моделі проходили однакові етапи:

- генерація набору даних;
- створення невизначеності;
- попередня обробка;
- навчання;
- тестування;
- оцінювання результатів.

Використання однакових умов дозволило виконати коректне порівняння алгоритмів.

Після завершення навчання система сформувала таблицю результатів.

Модель	Accuracy	Precision	Recall	F1-score
Decision Tree	0.303333	0.280033	0.303333	0.282941
Random Forest	0.320000	0.310372	0.320000	0.303761
Logistic Regression	0.310000	0.327442	0.310000	0.306610
Gaussian Naive Bayes	0.283333	0.295266	0.283333	0.287433

Проведений аналіз показав, що різні моделі демонструють різну якість прогнозування.

Прості алгоритми, такі як дерево рішень та наївний байєсівський класифікатор, забезпечують високу швидкість роботи, однак можуть втрачати точність при роботі з пошкодженими даними.

Ансамблеві методи, навпаки, демонструють більш стабільні результати.

У більшості випадків найкращі показники очікувано демонструє модель Random Forest, оскільки вона використовує сукупність дерев рішень та менш чутлива до шуму.

Якщо сказати простіше — одна модель може помилитися через дивні або пошкоджені дані, а велика група моделей працює стабільніше.

Висновки до розділу 3

У третьому розділі було реалізовано програмний модуль прогнозування дій супротивника та проведено експериментальні дослідження.

Було реалізовано механізм автоматичного створення синтетичного набору даних, моделювання невизначеності, попередньої обробки інформації та навчання моделей машинного навчання.

Також розроблено графічний інтерфейс користувача, який забезпечує зручну взаємодію із системою та дозволяє відображати результати прогнозування.

У результаті проведених експериментів виконано порівняння моделей машинного навчання та визначено найбільш ефективний алгоритм.

Результати дослідження показали, що використання ансамблевих методів дозволяє підвищити точність прогнозування та забезпечити стабільну роботу системи навіть за наявності неповної та невизначеної інформації.

ВИСНОВКИ

У кваліфікаційній роботі розглянуто задачу прогнозування дій супротивника в умовах неповної та невизначеної інформації із застосуванням моделей машинного навчання.

У ході виконання роботи було проведено аналіз сучасних підходів до прогнозування поведінки складних об'єктів та досліджено особливості використання алгоритмів машинного навчання для задач класифікації. Також виконано аналіз проблем, пов'язаних із неповнотою даних, шумом та невизначеністю інформації.

У процесі виконання роботи спроектовано архітектуру програмної системи, розроблено структуру синтетичного набору даних та визначено алгоритм функціонування програмного модуля.

Практична частина роботи включала реалізацію системи прогнозування засобами мови програмування Python. Для створення системи використовувалися бібліотеки машинного навчання, обробки даних, побудови графіків та створення графічного інтерфейсу.

У межах роботи реалізовано:

- механізм автоматичної генерації синтетичного набору даних;
- моделювання неповної та невизначеної інформації;
- алгоритми попередньої обробки даних;
- реалізацію моделей машинного навчання;
- механізм автоматичного вибору найкращої моделі;
- графічний інтерфейс користувача.

Під час експериментального дослідження було виконано навчання та тестування декількох моделей машинного навчання: дерева рішень, Random Forest, логістичної регресії та наївного байєсівського класифікатора.

У результаті проведених досліджень виконано порівняння моделей за основними показниками якості: accuracy, precision, recall та F1-score.

Результати експериментів показали, що використання ансамблевих методів машинного навчання дозволяє підвищити точність прогнозування та забезпечити більш стійку роботу системи в умовах пошкоджених даних.

Розроблена система успішно виконує прогнозування дій супротивника та може бути використана як основа для подальшого розвитку інтелектуальних систем підтримки прийняття рішень.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Pattern Recognition and Machine Learning / Christopher M. Bishop. – New York: Springer, 2006. – 738 p.
2. The Elements of Statistical Learning / T. Hastie, R. Tibshirani, J. Friedman. – Springer, 2009. – 745 p.
3. Hands-On Machine Learning with Scikit-Learn, Keras and TensorFlow / Aurélien Géron. – O'Reilly Media, 2022. – 851 p.
4. Artificial Intelligence: A Modern Approach / S. Russell, P. Norvig. – Pearson Education, 2021.
5. Data Mining: Practical Machine Learning Tools and Techniques / Ian Witten, Eibe Frank. – Morgan Kaufmann, 2016.
6. Python Official Documentation
7. NumPy Documentation
8. Pandas Documentation
9. Scikit-learn Documentation
10. Matplotlib Documentation
11. PyQt Documentation
12. Breiman L. Random Forests // Machine Learning. – 2001. – Vol.45(1). – P.5–32.
13. Quinlan J.R. Induction of Decision Trees // Machine Learning. – 1986. – Vol.1. – P.81–106.
14. Duda R., Hart P., Stork D. Pattern Classification. – Wiley-Interscience, 2012.
15. Han J., Kamber M., Pei J. Data Mining: Concepts and Techniques. – Morgan Kaufmann, 2011.
16. Murphy K. Machine Learning: A Probabilistic Perspective. – MIT Press, 2013.
17. Alpaydin E. Introduction to Machine Learning. – MIT Press, 2020.
18. James G., Witten D., Hastie T., Tibshirani R. An Introduction to Statistical Learning. – Springer, 2021.
19. Goodfellow I., Bengio Y., Courville A. Deep Learning. – MIT Press, 2016.

20. Bishop C. Neural Networks for Pattern Recognition. – Oxford University Press, 2005.
21. Vapnik V. Statistical Learning Theory. – Wiley, 1998.
22. Géron A. Machine Learning for Beginners and Advanced Users. – O'Reilly, 2022.
23. Chollet F. Deep Learning with Python. – Manning Publications, 2021.
24. Aggarwal C.C. Data Classification Algorithms and Applications. – CRC Press, 2014.
25. Raschka S. Python Machine Learning. – Packt Publishing, 2019.
26. Pedregosa F. et al. Scikit-learn: Machine Learning in Python // Journal of Machine Learning Research. – 2011.
27. McKinney W. Data Structures for Statistical Computing in Python // Proceedings of the Python Conference. – 2010.
28. Hunter J. Matplotlib: A 2D Graphics Environment // Computing in Science & Engineering. – 2007.
29. Pedrycz W., Chen S.-M. Machine Learning and Granular Computing. – Wiley, 2018.
30. Russell S., Norvig P. Artificial Intelligence and Intelligent Systems. – Pearson, 2021.
31. Закон України «Про інформацію».
32. Закон України «Про захист інформації в інформаційно-комунікаційних системах».
33. Методичні вказівки до виконання кваліфікаційної роботи.