

Національний університет "Львівська політехніка"

Інститут комп'ютерних наук та інформаційних технологій

Кафедра систем штучного інтелекту

Спеціальність 122 «Комп'ютерні науки»



ARTIFICIAL
INTELLIGENCE

ПОЯСНЮВАЛЬНА ЗАПИСКА

до магістерської кваліфікаційної роботи на тему:

«Застосування машинного навчання щодо оцінювання технічних характеристик важкої техніки для формування персоналізованих рекомендацій»

Студента(ки) групи

КНСШ-23

Лукаша О.В.

Керівник роботи

(підпис)

(Пелешишин О.П.)

Консультанти

(підпис)

(_____)

(підпис)

(_____)

Завідувач

кафедри СШІ

(Шаховська Н. Б.)

Національний університет "Львівська політехніка"
Інститут комп'ютерних наук та інформаційних технологій
Кафедра систем штучного інтелекту
Спеціальність 122 «Комп'ютерні науки»

ЗАТВЕРДЖУЮ:

Завідувач кафедри СШІ _____

“ ____ ” _____ 2023р.

ЗАВДАННЯ
НА МАГІСТЕРСЬКУ КВАЛІФІКАЦІЙНУ РОБОТУ СТУДЕНТОВІ

Лукашу Олексію Володимировичу

1. Тема роботи: «Застосування машинного навчання щодо оцінювання технічних аспектів важкої техніки для формування персоналізованих рекомендацій» затверджена наказом університету № 4420-4-06 від 13.10.2023
2. Термін подання закінченої роботи: __ 01.12.2023 _____
3. Вихідні дані до (проекту) роботи: вихідними даними виступатиме прогнозовані статистичні дані для персоналізованих рекомендацій щодо вибору техніки та ризиків її поломки за певними оставинами.
4. Зміст розрахунково-пояснювальної записки: (перелік питань, що належить розробити): аналітичний огляд літературних та інших джерел, системний аналіз, аналіз методів та засобів машинного навчання для програмної реалізації, практична реалізація.
5. Перелік графічного матеріалу (з точним зазначенням обов'язкових креслень): схема використання системи, схема роботи проаналізованих методів, візуалізація інтерфейсу користувача, представлення статистичних даних для персоналізованих рекомендацій.

6. Консультанти по роботі, із зазначенням розділів проекту, що стосується їх

Розділ	Консультант	Підпис, дата	
		Завдання видав	Завдання прийняв
Консультант з іноземної мови	Бойко Н. І.	 06.06.2023р	

7. Дата видачі завдання _____ 10.03.2023 _____

Керівник роботи

(підпис)

Завдання прийняв до виконання



(підпис)

КАЛЕНДАРНИЙ ПЛАН

№ п/п	Назва етапів дипломної роботи	Термін виконання етапів роботи	Примітка
1	Огляд та аналіз літературних джерел	10.03.2023 – 03.04.2023	
2	Постановка задачі дослідження та створення архітектури	03.04.2023 – 15.05.2023	
3	Розробка алгоритмічного та програмного забезпечення	15.05.2023 – 11.09.2023	
4	Експериментальні дослідження	11.09.2023 – 06.11.2023	
5	Оформлення пояснювальної записки	06.11.2023 – 30.11.2023	

Керівник роботи

(підпис)

Студент



(підпис)

АНОТАЦІЯ

Магістерська кваліфікаційна робота виконана студентом групи КНСШ-23 Лукашем Олексієм Володимировичем. Тема “Застосування машинного навчання щодо оцінювання технічних характеристик важкої техніки для формування персоналізованих рекомендацій”. Робота направлена на здобуття ступеня магістр за спеціальністю 122 «Комп’ютерні науки».

Метою дипломної роботи є формування персоналізованих рекомендацій щодо оцінювання важливих аспектів важкої техніки відповідно до сфери діяльності.

Об’єктом дослідження є важливі характеристики датчиків важкої техніки. У результаті виконання дипломної роботи було розроблено програмний продукт, який дозволяє отримати персоналізовані рекомендації щодо важкої техніки залежно від сфери та типу застосування техніки.

Предметом досліджень є методи та алгоритми обробки даних та машинного навчання для оцінювання технічних характеристик важкої техніки для формування персоналізованих рекомендацій.

Загальний обсяг роботи: 68 сторінок, 3 таблиці, 20 рисунків та 19 літературних джерел.

Ключові слова: персоналізовані рекомендації, машинне навчання, глибинна нейронна мережа, конволюційна нейронна мережа, Python scikit-learn.

ABSTRACT

Master's degree work of the student of the group CSAI-23 Lukash Oleksii Volodymyrovych. The topic is "Personalized recommendations system using machine learning to evaluate technical aspects of heavy equipmen". The work is aimed at obtaining a master's degree in 122 "Computer Science".

The purpose of the thesis is to form personalized recommendations for evaluating important aspects of heavy machinery in accordance with the field of activity.

The object of research is the important characteristics of heavy machinery sensors. As a result of the thesis, a software product was developed that allows you to get personalized recommendations for heavy machinery, depending on the scope and type of application of the equipment.

The subject of research is data processing and machine learning methods and algorithms.

Total amount of work: 68 pages, 3 tables, 20 figures and 19 references.

Keywords: personalized recommendations, machine learning, deep neural network, convolutional neural network, Python scikit-learn.

ЗМІСТ

АНОТАЦІЯ.....	4
ABSTRACT.....	5
ЗМІСТ.....	6
ВСТУП.....	7
1. АНАЛІТИЧНИЙ РОЗДІЛ.....	10
1.1. Постановка проблеми та її обґрунтування.....	10
1.2. Аналіз літературних джерел.....	11
1.2.1. Огляд літературних джерел.....	12
1.2.2. Критичний аналіз літературних джерел.....	15
1.3. Висновки до розділу.....	24
2. ДОСЛІДНИЦЬКИЙ РОЗДІЛ.....	26
2.1. Аналіз предметної області.....	26
2.1.1. Вхідні дані.....	27
2.1.2. Вихідні дані.....	28
2.2. Проектування системи.....	30
2.2.1. Обробка значень та структури даних.....	30
2.2.2. Вибір та аналіз моделей машинного навчання.....	33
2.2.3. Розбивання даних.....	37
2.2.4. Масштабування набору даних.....	40
2.2.5. CNN.....	42
2.2.6. LSTM.....	44
2.3. Інтерфейс програми.....	48
2.4. Висновки до розділу.....	48
3. РОЗДІЛ АПРОБАЦІЙ ТА РЕЗУЛЬТАТІВ.....	50
3.1. Засоби реалізації.....	50
3.2. Вимоги до технічного та програмного забезпечення.....	51
3.3. Опис програмної реалізації.....	52
3.3.1. Аналіз результатів.....	55
3.4. Керівництво користувача.....	58
3.5. Висновки до розділу.....	61
ВИСНОВКИ.....	63
СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ.....	66
Додаток А.....	68
Додаток Б.....	70

ВСТУП

Живучи у швидкоплинному та невизначеному світі, кожен член суспільства повинен приймати рішення щодо видів освіти та навчання, відкриття бізнесу, інвестиція чи будь яких інших дій, які забезпечують найкращу віддачу. Зазирнути в майбутнє неможливо в наш час, але попит на точно прогнозовану інформацію, рекомендації щодо прийняття рішень про потенційні події зростає.

У сучасному середовищі виникає потреба швидкого прийняття рішення. Наприклад, уряди хочуть зменшити тенденцію багатьох індексів, таких як безробіття, інфляція, промислове виробництво, а також спрогнозувати надходження від оподаткування, щоб сформувати ефективну політику. Менеджери з маркетингу хочуть збільшити попит на продукцію, обсяги продажів і зміни в уподобаннях споживачів, для цього потрібно прийняти ті чи інші рішення щодо поточної та майбутньої політики та, в більш загальному вигляді, сформулювати адекватне стратегічне планування. А формування рекомендацій щодо важкої техніки потрібно для людей, які займаються сільськогосподарською, будівельною, вантажною чи транспортною діяльністю. Адже саме техніка, її кількість, варіативність та технологічні можливості є основним предметом розвитку їх компаній.

Так як ринок важкої техніки є дуже великим та варіативним, тож це створює складності при пошуку техніки для потреб замовника. Для цього потрібно провести аналіз категорій техніки, оцінити всі важливі характеристики кожної з них відповідно до діяльності компанії. Ще одним важливим фактором є георозміщення та умови експлуатації самої техніки, адже відомо, що механізми, комплектуючі та навіть сам метал по різному зношується в залежності від навколишніх факторів, таких як постійна підвищена вологість, опади, низька температура. Тому є потреба в створенні рекомендаційної системи за допомогою методів машинного навчання. Застосування методів машинного навчання не обмежується особливостями однієї категорії техніки і є

гнучким. Для кращого результату потрібно налаштувати важливі параметри для дослідження та адаптуватись до будь яких змін даних.

Оскільки застосування важкої техніки у різних сферах життєдіяльності зростає відповідно до її кількості на первинному і вторинному ринку, то виникає нагальна потреба в аналізі та прийнятті рішень щодо спецтехніки. Обране дослідження є актуальними, що свідчить про великий попит на різного виду техніку та реальні труднощі з прийняттям рішень щодо її використання, експлуатування тощо.

Метою роботи є формування персоналізованих рекомендацій щодо оцінювання важливих аспектів важкої техніки відповідно до сфери діяльності.

Задачі, які потрібно виконати для досягнення поставленої цілі роботи:

- 1) Провести комплексний аналіз ринку важкої техніки та визначити найважливіші датчики для належного функціонування техніки.
- 2) Провести збір та критичну обробку даних, щоб підвищити точність прогнозування методів машинного навчання.
- 3) Провести аналіз моделей глибинного навчання, що включатиме створення архітектур моделей, які підходять для прогнозування даних.
- 4) Розробити персоналізовану систему рекомендацій для експлуатації важкої техніки зі зручним для користувача інтерфейсом.

Об'єктом дослідження являється ринок важкої техніки з оцінкою важливих характеристик.

Предметом дослідження є методи та засоби машинного навчання для обробки даних та формування рекомендаційних систем важкої техніки на ринку.

Наукова цінність роботи полягає у використанні методів машинного навчання для формування системи із персоналізованими рекомендаціями щодо важкої техніки.

Практична цінність роботи полягає у розробленні системи для прийняття рішень щодо експлуатації важкої техніки з оцінкою великої кількості параметрів.

Апробація результатів роботи. На основі проведених наукових досліджень було опубліковано наукову статтю по темі “Методика оцінки

вартості будівельної техніки на основі аналізу важливих характеристик з використанням методів машинного навчання” [1].

1. АНАЛІТИЧНИЙ РОЗДІЛ

1.1. Постановка проблеми та її обґрунтування

У сучасному світі, важка техніка відіграє важливу роль у більшості галузей промисловості. Забезпечення оптимальної та безперебійної роботи техніки є важливим, оскільки незаплановані збої та зупинки можуть призвести до значних фінансових витрат для ремонту або ж зовсім виведення із експлуатації самої техніки. Тому, розробка ефективної системи технічного обслуговування є важливою задачею для підприємств, що використовують важку техніку в своїй діяльності.

Одним з способів покращення ефективності системи технічного обслуговування є застосування машинного навчання для прогнозування можливих поломок. Це дозволить забезпечити належне технічне обслуговування та зменшить ризики виникнення незапланованих зупинок техніки. Для реалізації такої системи, необхідно врахувати безліч факторів, що можуть впливати на роботу важкої техніки. До таких факторів можна віднести стан техніки, середній час роботи машин, кліматичні умови, технічні характеристики тощо. Окрім цього, для формування персоналізованих рекомендацій необхідно враховувати особливості використання техніки на підприємстві та попередній досвід її експлуатації.

Отже, постановка проблеми полягає у тому, щоб розробити ефективну систему прогнозування поломок важкої техніки з використанням машинного навчання та створити персоналізовані рекомендації для її технічного обслуговування. Для покращення точності прогнозування необхідно використовувати різні методи машинного навчання та оптимізувати їх параметри з урахуванням конкретних особливостей використання важкої техніки. Крім того, для успішної реалізації системи необхідно мати якісні дані про стан машин та їх експлуатацію. Тому, необхідно використовувати сучасні датчики та інші пристрої, що забезпечують збір даних про стан техніки та її роботу.

Отже, розробка ефективної системи прогнозування поломок та

формування персоналізованих рекомендацій для технічного обслуговування важкої техніки з використанням машинного навчання є важливою задачею для підприємств, що використовують цю техніку в своїй діяльності.

1.2. Аналіз літературних джерел

Ми провели огляд статей, написаних англійською мовою, використовуючи такі бази даних: Scopus, Google Scholar, MDPI та інші дрібні платформи, дотримуючись рекомендацій Preferred Reporting Items for Systematic Reviews and Meta-analysis (PRISMA). На Рис.1.1 можна ознайомитися із схемою PRISMA, яка крок за кроком показує процес застосування критеріїв включення і виключення для генерації кінцевої кількості публікацій для мого запиту.

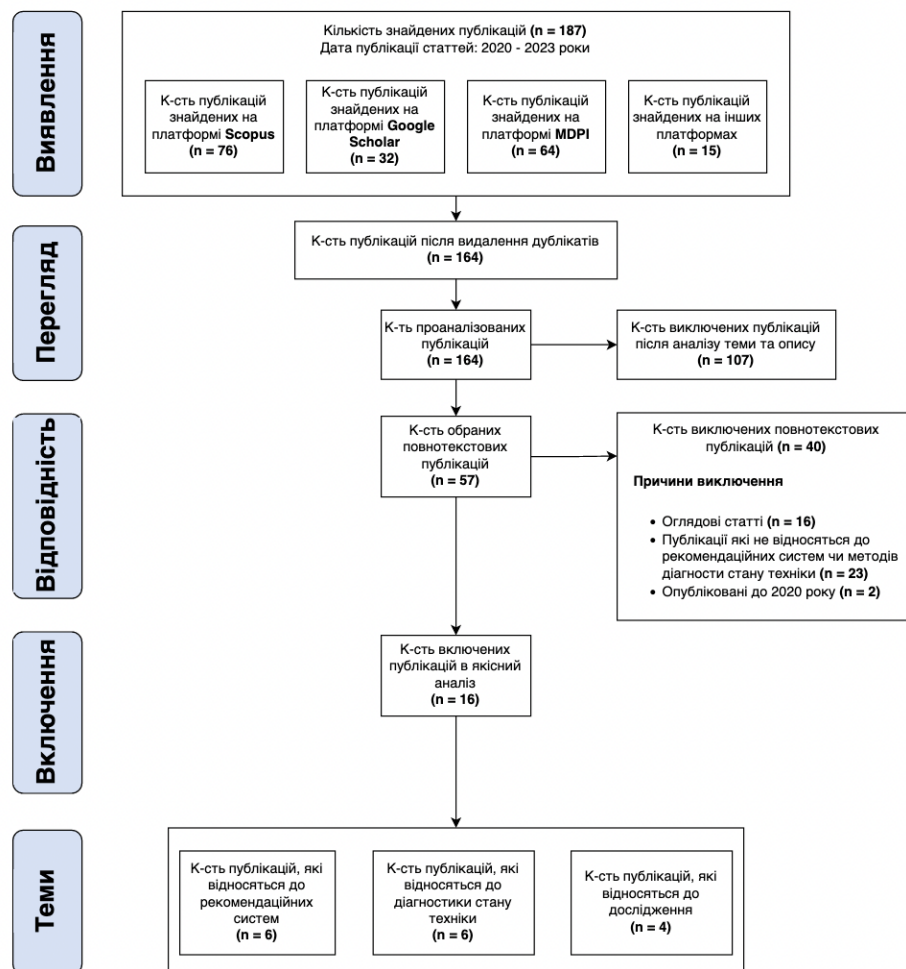


Рис. 1.1. Методологія виконання аналітичного огляду (PRISMA)

1.2.1. Огляд літературних джерел

На основі проведеного пошуку літературних джерел сформовано таблицю 1.1, яка представляє короткі відомості кожної із публікацій які будемо аналізувати в подальшому.

Таблиця 1.1. Огляд існуючих робіт

№	Назва	Інструментарій	Задача	Рік	Поси лання
1	Automatic learning algorithm for troubleshooting in hydraulic machinery	Machine Learning, TensorFlow	Розробити алгоритм автоматичного навчання для усунення несправностей у гідравлічних машинах	2023	[2]
2	Analyzing Reliability and Maintainability of Crawler Dozer BD155 Transmission Failure Using Markov Method and Total Productive Maintenance: A Novel Case Study for Improvement Productivity	ML, CNN, LSTM	Виявити основні причини несправностей трансмісії та запропонувати вдосконалення для підвищення надійності бульдозерів	2022	[3]
3	Evaluation of Agricultural Machinery Using Multi-Criteria Analysis Methods	Багатофакторні методи аналізу (AHP, TOPSIS, PROMETHEE)	Оцінити стійкість сільськогосподарськ ої техніки за допомогою методів багатофакторного аналізу.	2022	[4]
4	Location-aware personalized traveler recommender system (lapta) using collaborative filtering knn	ML, KNN, Python, Java	Розробити персоналізовані рекомендації щодо подорожей, використовуючи як місце розташування, так і налаштування користувача.	2021	[5]
5	Principal component analysis technique	PCA, Matlab	Розробити системи виявлення	2022	[6]

№	Назва	Інструментарій	Задача	Рік	Поси лання
	for early fault detection		несправностей для промислового обладнання		
6	Data-Driven Condition Monitoring of Mining Mobile Machinery in Non-Stationary Operations Using Wireless Accelerometer Sensor Modules	ML, Deep Learning, CNN, LSTM	Розробити систему моніторингу стану обладнання	2021	[7]
7	Designing predictive maintenance systems using decision tree-based machine learning techniques	ML, Decision Tree, Weka	Проектування систем прогнозного обслуговування з використанням методів машинного навчання на основі дерева рішень.	2020	[8]
8	Predictive Maintenance in Building Facilities: A Machine Learning-Based Approach	ML, CNN, Scikit-learn, Python	Розробити модель для прогнозування технічного обслуговування.	2021	[9]
9	An integrated deep learning-based approach for automobile maintenance prediction with GIS data	ML, LSTM, CNN	Розробити систему для прогнозування залишкового терміну служби автомобілів.	2021	[10]
10	Application of Machine Learning to a Medium Gaussian Support Vector Machine in the Diagnosis of Motor Bearing Faults	ML, SVM, MG-SVM	Розробити модель опорного вектора для діагностики стану підшипників.	2021	[11]
11	Predictive maintenance system	ML, Random Forest,	Розробити систему, яка зможе	2021	[12]

№	Назва	Інструментарій	Задача	Рік	Поси лання
	for production lines in manufacturing: A machine learning approach using IoT data in real-time	XGBoost, SVM, MQTT, Node-RED, Grafana	передбачити поломки на виробництві.		
12	SME 4.0: Machine learning framework for real-time machine health monitoring system	ML, Random Forest, Decision Tree, Neural Network	Розробити рекомендовану модель для прогнозування щодо технічного обслуговування	2021	[13]
13	A BiGRU method for remaining useful life prediction of machinery	Python, Keras, TensorFlow, Matlab	Підвищити точність прогнозування методу RUL шляхом використання різних підходів	2021	[14]
14	Machine learning models for predicting the residual value of heavy construction equipment: An evaluation of modified decision tree, LightGBM, and XGBoost regression	ML, Decision Tree, LightGBM, XGBoost	Розробити моделі прогнозування залишкової вартості, яка може використовуватися для прийняття обґрунтованих рішень щодо купівлі та продажу обладнання.	2020	[15]
15	Recommender systems and their ethical challenges	Recommendation systems	Дослідити етичні аспекти систем рекомендацій і надати рекомендації щодо розробки та регулювання цих систем.	2020	[16]
16	Deep Learning for Data-Driven Predictive Maintenance	Deep Learning, CNN, RNN	Дослідити потенціал методів глибокого навчання для прогнозного обслуговування обладнання.	2021	[17]

1.2.2. Критичний аналіз літературних джерел

Сучасний ринковий темп вимагає відповідної конкурентної переваги. Тобто, якщо прагнути зайняти міцну, прибуткову позицію серед бізнес-конкурентів у своїй ніші. Для цього слід використовувати відповідні технології, які допоможуть залишатися на крок попереду від своїх конкурентів. На процес прогнозування технічного обслуговування та пошук збоїв в обладнанні впливають технології, що прискорюють обробку даних. В даній роботі на пришвидшення обробки даних можна вплинути за допомогою методів машинного навчання.

Метою дослідження [2] є розробка алгоритму, який може знаходити проблеми в системах гідравлічних машин. У статті детально розглядається побудова алгоритму та його специфіка. Результати дослідження можуть мати значний вплив на галузь гідротехнічного машинобудування, покращуючи технічне обслуговування машин та зменшуючи час простою на ремонті. Висновки статті також можуть бути корисними у контексті мого дослідження. Модель була навчена на типових несправностях гідравлічного обладнання та відповідних симптомах, щоб визначити найбільш ймовірну причину несправності. Дослідники протестували алгоритм на змодельованій гідравлічній системі з кількома несправностями, алгоритм точно виявив усі несправності, продемонструвавши його потенціал для практичного застосування. Автори також порівняли свій алгоритм з існуючими методами діагностики, такими як експертні системи, засновані на правилах, і виявили, що підхід на основі машинного навчання є більш точним і швидшим. Таким чином, дослідження представляє новий і перспективний підхід до усунення несправностей гідравлічних систем з потенційним застосуванням у різних галузях промисловості.

У статті [3] розповідається про використання методу Маркова та повного продуктивного технічного обслуговування для виявлення основних причин несправностей трансмісії. У дослідження також було запропоновано способи їх усунення для підвищення надійності та ремонтпридатності гусеничного бульдозера. Автори зібрали дані про технічне обслуговування бульдозера та

використали метод Маркова для аналізу надійності та ремонтпридатності системи трансмісії. Вони також запровадили практику повного продуктивного технічного обслуговування, щоб покращити якість обслуговування та скоротити час простою для ремонту. Результати дослідження показали, що ці методи допомогли виявити основні причини несправностей та запропонувати рішення для вирішення цих проблем. Все це допоможе підвищити надійність та ремонтпридатність гусеничного бульдозера, що може призвести до підвищення продуктивності використання техніки. Ця стаття може бути корисною для людей, які працюють у будівельній та гірничодобувній промисловості, де використовуються гусеничні бульдозери.

Стаття [4] досліджує, наскільки стійкі сільськогосподарські машини за допомогою методів багатофакторного аналізу. У статті наведено кілька методів для оцінки стійкості сільськогосподарської техніки: метод аналізу ієрархій, метод прийняття рішень за схожістю з ідеальним вибором і метод рейтингової оцінки. Автори дослідили різні критерії стійкості, такі як енергоефективність, вплив на навколишнє середовище, економічну доцільність і соціальну прийнятність, використовуючи дані з різних джерел. Вони використали вищезгадані методи, щоб визначити вагу критеріїв стійкості, спробували розміщувати машини за шкалою стійкості, та визначити найбільш стійкі машини на основі їх рейтингу. Дана праця є корисною для запропонованої теми, адже тут описуються методи багатофакторного аналізу техніки, які будуть застосовувати на практичній частині дослідницької роботи. Дослідження також дозволяє краще зрозуміти, наскільки стійкі сільськогосподарські машини за допомогою різних методів та критеріїв.

У статті [5], опублікованій у 2021 році в журналах Computers, Materials and Continua, пропонується інноваційна система рекомендацій для подорожей, яка враховує як уподобання користувача, так і дані про місцезнаходження. Використовуючи алгоритм К-найближчого сусіда (KNN), дослідження має на меті запропонувати індивідуальні та точні пропозиції щодо подорожей на основі історії подорожей користувача, його оцінок та попереднього вибору. Колаборативна фільтрація є основним методом, який використовується в цій

системі, що дозволяє моделі збирати дані з різних джерел для підвищення точності рекомендацій. Дослідження порівнює систему з іншими системами, такими як матрична факторизація та фільтрація на основі контенту, і результати показують, що створена система перевершує їх з точки зору точності та релевантності. Загалом, “персоналізована система мандрівникам з урахуванням місцезнаходження” є цінним інструментом для мандрівників, які шукають персоналізовані рекомендації, що відповідають їхнім уподобанням та інтересам. Вона без ніяких труднощів поєднує персональні дані та інформацію про місцезнаходження, щоб генерувати індивідуальні рекомендації для користувачів.

Стаття [6] присвячена системі, яка використовує аналіз головних компонент (РСА) для раннього виявлення несправностей у промисловому обладнанні. РСА - це підхід до аналізу даних, який може розпізнавати приховані закономірності та взаємозв'язки у складних наборах даних. У дослідженні було проведено РСА-аналіз сигналів вібрації промислового обладнання та використано ці результати для побудови системи раннього виявлення несправностей. Основним інструментом, який використовували автори для розробки алгоритму РСА та системи раннього виявлення несправностей, був Matlab. Для аналізу даних та оцінки продуктивності системи також використовувалися статистичні інструменти, такі як середнє значення, стандартне відхилення та дисперсія. Експериментальні дані були зібрані з коробки передач за різних умов експлуатації та несправностей. Дані були попередньо оброблені і відфільтровані перед використанням в алгоритмі РСА для вилучення основних компонентів. Потім автори розробили модель виявлення несправностей на основі цих компонентів з використанням порогового підходу. У дослідженні було оцінено продуктивність системи раннього виявлення несправностей за допомогою таких показників, як чутливість, специфічність, точність і крива робочої характеристики приймача (ROC). Було виявлено, що система на основі РСА є надійною і ефективною у виявленні несправностей на ранніх стадіях. Загалом, автори використовували РСА, статистичні показники та Matlab для розробки системи раннього

виявлення несправностей для промислового обладнання. Це дослідження висвітлює потенціал методів на основі РСА у своєчасному виявленні несправностей, відкриваючи шляхи для майбутніх досліджень.

Метою дослідження [7] авторів була розробка системи моніторингу стану мобільного гірничодобувного обладнання за допомогою відстеження даних. Вони досягли цієї мети, використовуючи бездротові модулі датчиків для збору даних про вібрацію під час руху обладнання. Для вивчення даних і визначення стану обладнання автори використовували передові методи машинного навчання, включаючи алгоритми моделей LSTM і CNN. Щоб підготувати дані для моделей машинного навчання, автори використовували методи попередньої обробки даних: нормалізація, стандартизація та вилучення ознак. Для підготовки даних для моделей машинного навчання були використані методи нормалізації, вилучення ознак і попередньої обробки даних для покращення якості. Для оцінки ефективності системи використовували середньоквадратичну похибку (RMSE) та метрики матриці помилок (Confusion matrix). Таким чином, дослідження продемонструвало, що бездротові сенсорні модулі та технології машинного навчання мають потенціал для створення систем, керованих даними, здатних контролювати стан мобільного гірничодобувного обладнання в нестационарних умовах.

У роботі [8] представлено дослідження, яке вивчає застосування методів машинного навчання на основі дерев рішень для проектування систем превентивного технічного обслуговування. У дослідженні були використані дані з виробничого заводу в реальному контексті. Автори запропонували п'ятиетапну структуру для проектування систем предиктивного обслуговування: (1) збір даних, (2) попередня обробка даних, (3) вилучення ознак, (4) побудова моделі та (5) оцінка моделі. Алгоритм дерева рішень був використаний як основний метод машинного навчання для побудови моделей, а кілька оціночних показників, таких як точність, достовірність, пригадування та оцінка F1, були використані для вимірювання продуктивності моделі. Результати показали, що алгоритм дерева рішень точно передбачив вихід обладнання з ладу, виділивши найбільш важливі функції для подальшого прогнозованого обслуговування. Дослідники

дійшли висновку, що методи машинного навчання на основі дерев рішень можна вважати надійним інструментом для розробки систем превентивного технічного обслуговування.

У статті [9] запропоновано використовувати машинне навчання для прогнозованого обслуговування будівель. Модель прогнозує несправності обладнання, використовуючи різні алгоритми (дерева рішень, метод опорних векторів та k-найближчих сусідів). Крім того, у дослідженні описано процес кореляційного аналізу та відбору ознак, які використовувалися для визначення основних характеристик, необхідних для побудови моделі. Модель навчається та оцінюється за допомогою очищеного набору даних, отриманих з датчиків будівлі, який був розділений на навчальний та тестовий набори. Для оцінки її ефективності використовуються різні метрики точності. Цей підхід має потенціал для зменшення витрат на технічне обслуговування і підвищення надійності та ефективності обладнання, що робить його життєздатним рішенням, саме через цей критерій, було обрано цю статтю для аналізу літератури для моєї предметної області.

У публікації [10] представлено підхід до прогнозування технічного обслуговування транспортних засобів, який використовує дані географічних інформаційних систем та глибоке навчання. Автори пропонують структуру, яка інтегрує згорткові нейронні мережі (CNN) і мережі з довгою короткочасною пам'яттю (LSTM) для прогнозування залишкового терміну експлуатації (RUL) транспортних засобів. У дослідженні використовується набір даних із записами про технічне обслуговування 399 автомобілів за три роки, а також відповідні ГІС-дані. Автори попередньо обробляють дані, виділяють і відбирають ознаки та навчають модель CNN-LSTM. Для оцінки якості моделі, вони використовують середню абсолютну похибку та середньоквадратичну похибку. Результати демонструють, що запропонований підхід перевершує традиційні алгоритми машинного навчання в прогнозуванні для типових у дослідженні даних.

У статті автора Ши-Лін Лін [11], опублікованій у 2021 році, представлено застосування машинного навчання для діагностики несправностей підшипників

двигуна. Автор пропонує модель опорного вектора Гаусса (MG-SVM), яка використовує ознаки, отримані з вібраційних сигналів, для діагностики стану підшипників. В ході дослідження було зібрано сигнали вібрації зі стенду для випробування підшипників у двигуну та переформатовано у числовий формат, який буде сприймати модель. Модель MG-SVM була навчена з використанням цих функцій для класифікації сигналів вібрації як функціональних або дисфункціональних. Ефективність моделі оцінювали за кількома показниками, включаючи точність, чутливість, специфічність і площу кривої робочої характеристики приймача. Результати показали, що запропонована модель MG-SVM досягла високого рівня точності і перевершила інші моделі машинного навчання, які використовувалися для цієї ж задачі. Дослідження продемонструвало ефективність запропонованої моделі в діагностиці несправностей підшипників двигуна і показало її потенціал для застосування в реальних сценаріях. Цю статтю було обрано, щоб допомогти реалізувати алгоритм отримання даних з різних датчиків, тож ця публікація цілком закриває одну із моїх потреб у дослідженні.

У роботі [12] автори пропонують використовувати алгоритми машинного навчання для створення системи прогнозованого обслуговування виробничих ліній, яка може заздалегідь виявляти потенційні збої, знижуючи витрати на обслуговування і підвищуючи продуктивність. Для цього автори використовували різні інструменти і методи, включаючи збір даних з датчиків, розташованих на виробничих лініях, і аналіз даних в режимі реального часу. Використовуючи алгоритми машинного навчання, такі як Random Forest, XGBoost і Support Vector Machine, автори змогли розробити точні прогнозні моделі для прогнозування потреб у технічному обслуговуванні. Крім того, автори використовували MQTT (Message Queuing Telemetry Transport), Node-RED і Grafana, щоб полегшити передачу даних, візуалізацію і прийняття рішень в режимі реального часу, тим самим підвищивши ефективність запропонованої системи, це і є один із критеріїв, обрання даної роботи для аналізу. Впровадження системи дало багатообіцяючі результати, які свідчать про те, що вона ефективно прогнозує та запобігає виходу обладнання з ладу,

підвищуючи загальну продуктивність. Підхід, заснований на машинному навчанні, використовував дані в режимі реального часу для прогнозування потреб у технічному обслуговуванні, що дозволило мінімізувати простой виробничої лінії. Автори отримали дані за допомогою декількох датчиків, включаючи датчики температури, тиску і вібрації, які вони проаналізували за допомогою дерев рішень і випадкових лісів. Запропонована ними система дала точні результати, прогнозуючи несправності обладнання. Дослідження підкреслює потенціал машинного навчання у прогнозуванні технічного обслуговування для виробничих галузей.

У дослідженні [13] детально описується гнучкий і адаптований фреймворк машинного навчання для моніторингу стану обладнання в режимі реального часу на малих і середніх підприємствах у виробничій галузі. Автори пропонують чотириетапну структуру, що включає (1) збір даних, (2) вилучення ознак, (3) вибір ознак і (4) прогнозне моделювання з використанням різних алгоритмів машинного навчання. Оцінюючи ефективність системи, автори дійшли висновку, що фреймворк виявляє і класифікує різні типи несправностей верстата (стукіт і знос інструменту) з високою точністю. Ця стаття цінна тим, що автори надають детальний опис реалізації концепції, що робить її доступною як для практиків, так і для дослідників. Проте одним з недоліків статті є відсутність порівняння з існуючими системами машинного навчання або традиційними статистичними методами, чим ускладнює оцінку переваги запропонованої системи над іншими підходами. Запропонована система має потенціал для підвищення ефективності та продуктивності виробничого процесу, скорочення витрат на технічне обслуговування, а також має практичне застосування.

У статті [14], опублікованій у 2021 році, пропонується новий метод прогнозування залишкового ресурсу машин. Автори зосереджуються на використанні двонаправлених блоків повторення (BiGRU) для моделювання послідовних даних і прогнозування залишкового ресурсу. Вони застосували запропонований метод до двох прикладів: (1) прогнозування залишкового ресурсу імітованого турбовентиляторного двигуна і (2) прогнозування

залишкового ресурсу реальної коробки передач. Результати показують, що запропонований метод BiGRU перевершує інші існуючі методи з точки зору точності прогнозування RUL. Метод BiGRU, був реалізований за допомогою мови програмування Python та різних її бібліотек (TensorFlow і Keras). Вони також використовують програмне забезпечення MATLAB для порівняння з іншими методами прогнозування RUL. Автори попередньо обробляють дані, масштабуючи і нормалізуючи їх, і застосовують метод згладжування (підхід ковзного вікна) для підготовки даних для моделі BiGRU. Для оцінки ефективності, вони використовують різні метрики (середню абсолютну похибку і середньоквадратичну похибку), щоб порівняти точність свого методу з іншими методами. В цілому, метод BiGRU для машинного прогнозування RUL, запропонований в цій статті, є інноваційним та ефективним підходом до прогнозного технічного обслуговування.

У 2021 році було опубліковано статтю [15]. Метою дослідження було створення надійної моделі для прогнозування залишкового ресурсу важкого будівельного обладнання, яку компанії, що надають обладнання в оренду, могли б використовувати їх для прийняття обґрунтованих рішень щодо купівлі та продажу самого обладнання. Для цього автори зібрали дані про важку будівельну техніку з різних джерел, провели попередню обробку даних, щоб усунути будь-які відсутні або аномальні дані. Потім вони оцінили ефективність трьох моделей машинного навчання, за декількома метриками, середня абсолютна відсоткова похибка (MAPE), середньоквадратична похибка (RMSE) та коефіцієнт детермінації (R^2). Результати показали, що з трьох протестованих моделей регресійна модель XGBoost випередила інші моделі, зафіксувавши найнижчі значення RMSE, MAPE та найвищі значення R^2 . Тож, автори дійшли висновку, що регресійна модель XGBoost є надійним засобом прогнозування залишкової вартості важкої будівельної техніки, що доводить її цінність для компаній, які займаються орендою техніки.

Маріаросарія Таддео та Лучано Флоріді [16], досліджують етичні проблеми, що виникають у зв'язку з використанням рекомендаційних систем. Рекомендаційні системи використовують алгоритми машинного навчання для

надання індивідуальних рекомендацій користувачам на основі їхньої попередньої поведінки, інтересів та уподобань. Однак автори стверджують, що ці системи створюють етичні проблеми, пов'язані з конфіденційністю, автономією, справедливістю та підзвітністю. Для боротьби з цими дилемами автори пропонують концепцію, яка інтегрує етичні міркування на всіх етапах проектування, розробки та впровадження рекомендаційних систем. Система складається з чотирьох ключових аспектів: (1) конфіденційність і захист даних, (2) прозорість і пояснення, (3) справедливість і недискримінація, а також (4) підзвітність і управління, що вимагає міждисциплінарного підходу, який передбачає залучення комп'ютерників, інженерів, соціологів, юристів і етиків. Автори також обговорюють кілька етичних дилем у рекомендаційних системах, включаючи напругу між персоналізацією і конфіденційністю та ризик виникнення "бульбашок фільтрів". Вони рекомендують, щоб етичні міркування були вплетені в алгоритми для пом'якшення цих проблем. На завершення стаття підкреслює необхідність етичних рекомендаційних систем для захисту конфіденційності, автономії, справедливості та підзвітності користувачів. Дотримуючись запропонованої концепції, алгоритми можуть бути розроблені таким чином, щоб інтегрувати етичні міркування і гарантувати, що вони сприятимуть здоровій соціальній взаємодії та інноваціям, не використовуючи при цьому персональні дані.

Проаналізувавши статтю [17], стало очевидно, що Шиза Муштак і Гіа Уддін приділяють значну увагу тому, як впровадження прогнозованого технічного обслуговування може підвищити надійність обладнання та зменшити витрати на обслуговування. Вони пропонують використовувати технологію глибокого навчання як надійний засіб прогнозування потенційних збоїв обладнання та машин, тим самим запобігаючи неочікуваним простоям і підвищуючи рівень продуктивності. У своєму дослідженні прогнозованого обслуговування автори аналізують різні набори даних, включаючи дані датчиків і дані технічного обслуговування, а також розглядають пов'язані з цим функції. Вони також описують проблеми із якими стикнулися під час роботи з цими

даними. Однієї із дискусій дослідження є оцінка методів глибокого навчання, таких як рекурентні нейронні мережі та згорткові нейронні мережі.

1.3. Висновки до розділу

В ході дослідження за темою було сформовано пошукові запити до наукометричних баз – Scopus, Google Scholar та MDPI. За допомогою методології виконання аналітичного огляду (PRISMA) було відібрано 16 релевантних статей та сформовано таблицю 1.1 із короткими відомостями про кожну з них.

На одному з етапів даної роботи було проведено критичний огляд літератури, у яких було розглянуто різні методи машинного навчання, алгоритми прогнозування технічного обслуговування, та створення рекомендаційних систем. Так як застосування машинного навчання дає можливість прогнозувати поломки техніки та забезпечувати її надійність, а персоналізовані рекомендації щодо технічного обслуговування дозволяють зменшити витрати на ремонт та збільшити термін експлуатації техніки.

Аналізуючи дослідження також варто згадати труднощі із якими можна зіткнутися у процесі створення рекомендаційної системи для прогнозованого технічного обслуговування. Серед таких складностей можна виділити:

- наявність великих обсягів даних щодо технічного стану (переважно значення зібрані із датчиків) обладнання;
- система повинна бути адаптована до різних типів важкої техніки та умов експлуатації.

Підбиваючи підсумки, можна зробити висновок, що пошук ефективних методів прогнозування поломок важкої техніки та формування персоналізованих рекомендацій є актуальним завданням для багатьох галузей виробництва, які стикаються з проблемами ефективного технічного обслуговування.

Заключно можна сказати, що дослідження застосування машинного навчання для прогнозування технічного обслуговування важкої техніки та формування персоналізованих рекомендацій може стати важливим внеском у

розвиток індустріальної автоматизації та мати великий практичний потенціал для застосування в реальному виробничому середовищі. Після завершення, можна проводити подальші дослідження для покращення алгоритмів та моделей машинного навчання, що дозволить ще більш ефективно використовувати дані методи в галузі промислового виробництва та технічного обслуговування.

2. ДОСЛІДНИЦЬКИЙ РОЗДІЛ

2.1. Аналіз предметної області

Для вирішення поставлених в роботі завдань буде використовуватися підхід прикладні методи на прикладі прогнозування часових рядів.

Багато проблем прогнозування включають часовий компонент і, таким чином, вимагають екстраполяції даних часових рядів, або прогнозування часових рядів. Прогнозування часових рядів також є важливою сферою машинного навчання (ML) і може розглядатися як проблема керованого навчання. До нього можна застосувати такі методи ML, як регресія, нейронні мережі, машини опорних векторів, випадкові ліси та XGBoost. Прогнозування передбачає використання моделей, пристосованих до історичних даних, для передбачення майбутніх спостережень.

Прогнозування часових рядів означає передбачення або передбачення майбутнього значення протягом певного періоду часу. Воно передбачає розробку моделей на основі попередніх даних і застосування їх для спостережень і прийняття майбутніх стратегічних рішень.

Майбутнє прогнозується або оцінюється на основі того, що вже відбулося. Часовий ряд додає залежність часового порядку між спостереженнями. Ця залежність є одночасно і обмеженням, і структурою, яка забезпечує джерело додаткової інформації. Перш ніж ми обговоримо методи прогнозування часових рядів, давайте визначимо, що таке прогнозування часових рядів.

Прогнозування часових рядів - це метод передбачення подій через певну послідовність у часі. Він передбачає майбутні події, аналізуючи тенденції минулого, виходячи з припущення, що майбутні тенденції будуть подібними до історичних. Він використовується в багатьох галузях науки в різних додатках, зокрема: астрономія, бізнес-планування, контрольно-вимірювальна техніка, прогнозування землетрусів, економетрика, математичні фінанси, розпізнавання образів, розподіл ресурсів, обробка сигналів, статистика, прогнозування погоди.

2.1.1. Вхідні дані

Вхідними даними в основній частині реалізованої системи сформовані та

структуровані набори показів датчиків з основних вузлів важкої техніки для тренування моделей машинного навчання. Всі покази датчиків являють собою числові значення з різним розкидом величин (див. Рис. 2.1 та 2.2).

	Date	Engine Temperature	Engine RPM	Fuel Consumption	Hydraulic Flow Rate	Hydraulic Pressure	Hydraulic Temperature
0	2023-04-01 00:00:00	29.664949	29.010285	30.397275	31.009183	29.621066	29.721267
1	2023-04-01 00:01:00	29.664949	29.010285	30.397275	31.009183	29.621066	29.721267
2	2023-04-01 00:02:00	29.416357	29.170713	30.397275	31.067310	29.831835	28.588632
3	2023-04-01 00:03:00	29.605749	29.010285	30.372483	31.067309	29.329233	29.926988
4	2023-04-01 00:04:00	29.428197	29.037021	30.397275	31.067309	29.718342	29.771497
...	...	...	...	...	...	...	...
114045	2023-06-19 04:45:00	28.848156	29.384612	28.884855	28.887657	29.496769	25.585060
114046	2023-06-19 04:46:00	28.931012	29.384612	28.884852	28.887657	29.401384	26.299007
114047	2023-06-19 04:47:00	28.931012	29.384612	28.884855	28.887657	29.401382	26.299007
114048	2023-06-19 04:48:00	28.907344	29.384612	28.860057	28.887657	29.372467	25.987065
114049	2023-06-19 04:49:00	28.931012	29.384612	28.860057	28.887657	29.469747	25.864708

114050 rows × 7 columns

Рис. 2.1. Приклад показів даних із датчиків

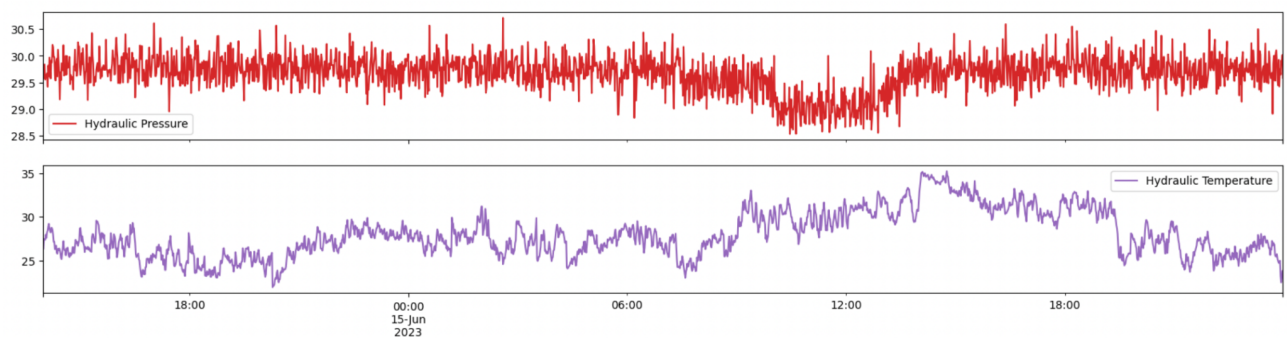


Рис. 2.2. Візуальне представлення даних датчиків гідравлічного тиску та температури на короткому проміжку числового ряду

Ще одним етапом системи є створення персоналізованих рекомендацій стосовно експлуатації техніки. Вхідними даними для цього етапу слід вважати важку техніку, характеристики якої та персональні дані вже внесені в базу даних, а також вид діяльності та складність виконання роботи, що буде впливати на покази з датчиків важкої техніки.

Також було прийнято рішення створити зручний користувацький інтерфейс для роботи із системою (див. Рис. 2.3).

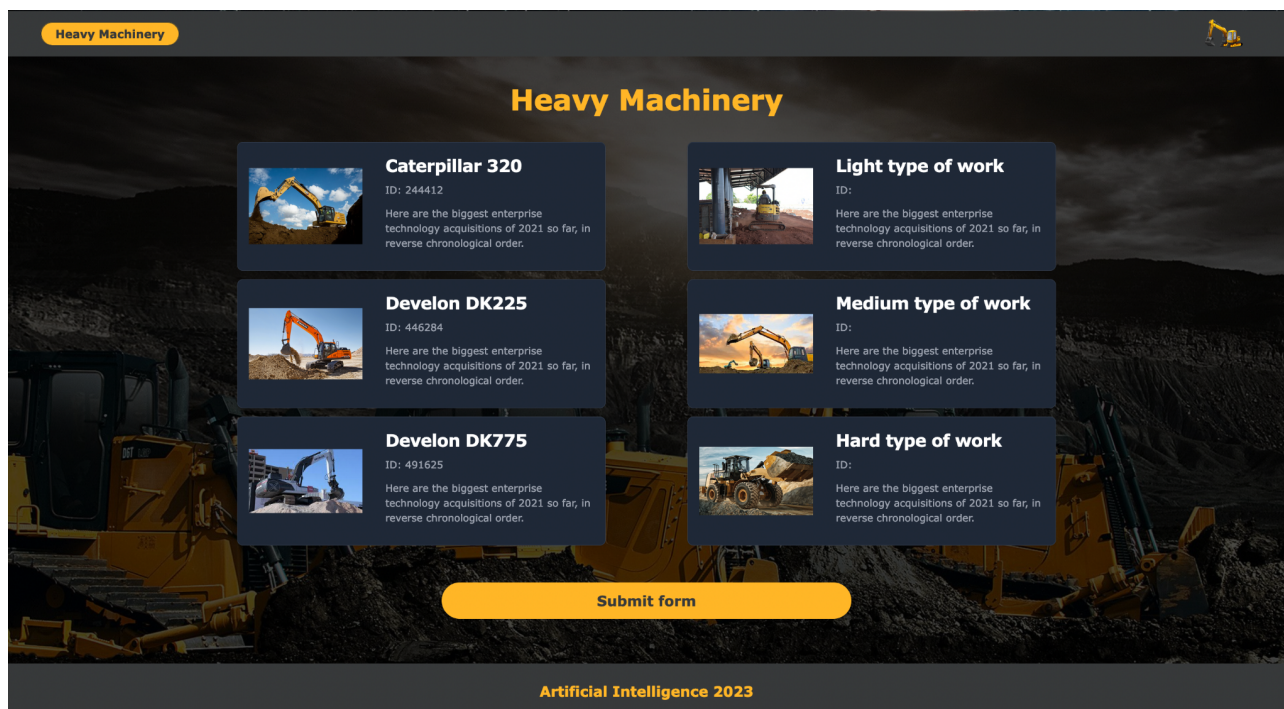


Рис. 2.3 Графічний користувацький інтерфейс розробленої програми

Нам потрібно обрати техніку із доступних та вид складності діяльності, на яку планується відправити техніку. Ці параметри відправляються на сервер та обробляються раніше натренованими моделями машинного навчання.

2.1.2. Вихідні дані

Частина системи, яка відповідає за пошук буде повертати великі обсяги даних із різних датчиків системи важкої техніки.

Далі модуль відповідальний за обробку, повертатиме нам єдиний потік даних, який міститиме попередні значення показів разом із новими значеннями важливих характеристик транспортного засобу.

Наступною і основною частиною системи є побудова та навчання моделі машинного навчання для отримання персоналізованих рекомендацій щодо правильності експлуатування важкої техніки. Вихідними даними є сама модель та результати оцінки моделі.

Завершальною частиною системи є інтерфейс. Вихідними даними виступатиме персоналізовані рекомендації щодо використання транспортного засобу та ймовірність виходу з ладу при певних умовах експлуатації техніки (див. Рис. 2.4).



Result



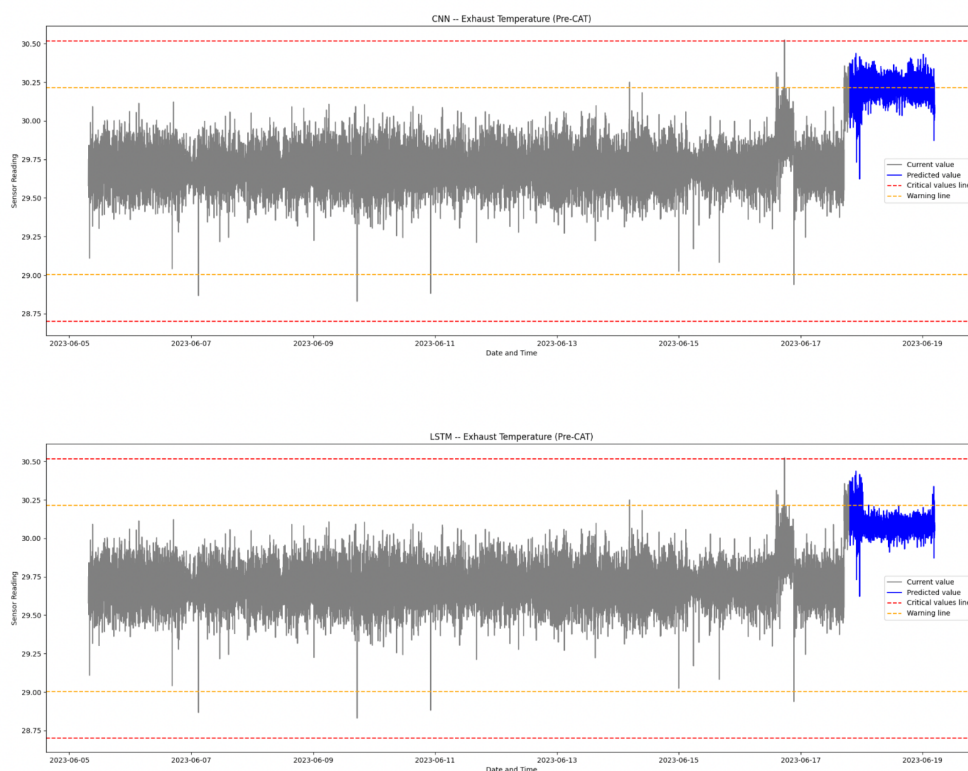
Develon DK225

ID: 446284

The selected technique for this particular task is ill-suited and raises concerns about its performance. The technique in question lacks the necessary robustness and resilience required for the task, making it vulnerable to malfunctions and damage. The equipment and processes associated with this technique are known to be susceptible to stress and wear and tear when subjected to the demands of the given work. The risk of breakdowns is high, and repairs may be necessary, adding both time and costs to the project. You can potentially experience breakdowns identified by the following sensors:

Voltage Supply, Exhaust Temperature (Pre-CAT), Oil Pressure, Transmission Oil Pressure, Bucket Position, Bucket Angle.

Pay attention to checking the sensors and the unit as a whole.



Artificial Intelligence 2023

Рис. 2.4 Користувачський інтерфейс із рекомендаціями щодо обраної техніки та ступеня складності робіт

На графіках зображено оригінальні значення із датчиків та прогнозовані значення відповідно до кожної моделі нейронної мережі. На графіках також присутні межі мало- та сильно-критичних показів, вони обраховані за допомогою середнього відхилення усіх оригінальних значень датчиків. Такий підхід було обрано з огляду на практичність, в реальності ці межі

встановлюються виробниками самих датчиків та залежать від характеристик кожної моделі техніки індивідуально.

За допомогою цих меж, система обраховує кількість таких значень і якщо вона перевищує 10% всіх прогнозованих значень, то такий датчик схильний до збоїв та потребує планового ремонту чи заміни, про що і сигналізує система.

2.2. Проектування системи

Отримані дані показів із основних вузлів важкої техніки, потрібно привести до одного вигляду та структури. Для цього ми будемо використовувати окремий модуль для обробки даних та структурування їх відповідно до категорії, типу та моделі важкої техніки.

У цій роботі було використано чималу кількість різних датчиків, значення з яких представлені в часових рядах. Дані - це регулярні покази з датчиків, зібрані з 46 основних параметрів важкої техніки: датчики температури, датчики тиску, датчик потоку охолодження, кількість викидів, тощо. Дані охоплюються покази за 3 місяці із похвилинним записом значень. Покази із датчиків будуть прогнозуватися індивідуально з використанням історичних значень за два дні спостережуваних рядів у цьому експерименті.

2.2.1. Обробка значень та структури даних

Попередня обробка даних у машинному навчанні є важливим кроком, який допомагає підвищити якість даних, отримуючи лише релевантну інформацію. Під попередньою обробкою даних в машинному навчанні розуміють підготовку (очищення та організацію) необроблених даних, щоб зробити їх придатними для побудови та навчання різних моделей машинного навчання.

Попередня обробка даних у машинному навчанні складається з декількох етапів:

- Пошук і видалення дублікатів даних.
- Видалення аномальних даних (викидів).

Критично важливим етапом попередньої обробки даних є видалення так званих викидів. Викиди - це специфічні об'єкти, які сильно відрізняються від

генеральної сукупності. Вони мають характеристики, які відрізняються від більшості інших об'єктів у наборі даних (див. Рис. 2.5).

Викиди бувають двох типів:

- Одновимірний викид - це точка даних, яка складається з екстремального значення однієї змінної.
- Багатовимірний викид - це комбінація незвичайних значень принаймні для двох змінних.

Наявність викидів у наборі даних може негативно вплинути на якість навчання моделі машинного навчання. Тому основним етапом попередньої обробки даних є видалення викидів.

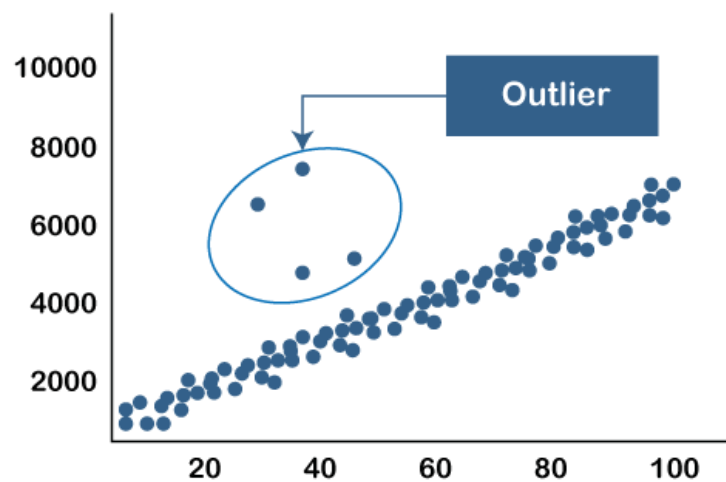


Рис. 2.5. Виявлення викидів

Для виявлення викидів використовують два популярних методів[1]:

- Метод середньоквадратичного відхилення (Standard deviation method);
- Міжквартильний інтервал або метод міжквартильного інтервалу (IQR: Interquartile Range or Interquartile Range Method).

Метод середньоквадратичного відхилення використовується тільки при умові нормального розподілу вхідних даних. В нормальному розподілі, дані розподілені симетрично. Більшість значень зосереджені навколо центральної області, зменшуючись у міру віддалення від центру. Середньоквадратичне відхилення показує, на скільки позицій в середньому розташовані дані від центру розподілу.

Нормальний розподіл містить два важливих параметри - середнє значення і стандартне відхилення. У нормальному розподілі ці два параметри можна використовувати для оцінки нетипових даних у вибірці.

Стандартне відхилення і середнє значення можуть вказувати на положення значень у нормальному розподілі. Емпіричне правило, або правило 68-95-99,7, стверджує, що 68% даних знаходяться в межах одного стандартного відхилення, двох стандартних відхилень - 95%, а трьох стандартних відхилень - 99,7% (Рис. 2.6).

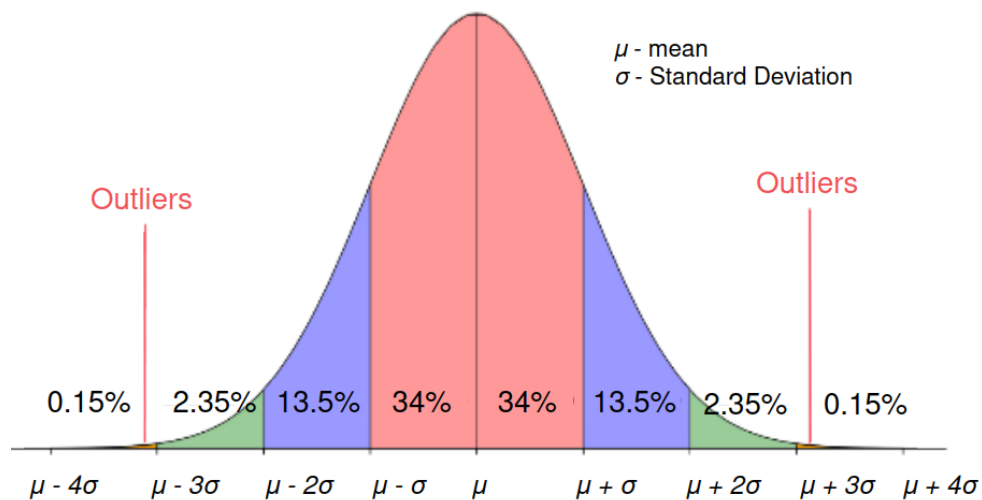


Рис. 2.6. Метод середньоквадратичного відхилення

Дані, що виходять за межі трьох стандартних відхилень, також є частиною генеральної сукупності, але це нетипові або рідкісні випадки. Тому, як правило, всі дані, що виходять за межі трьох стандартних відхилень, вважаються викидами і повинні бути вилучені з сукупності. Це значення може змінюватися залежно від розміру набору даних: якщо набір даних занадто великий, дані можуть виходити за межі чотирьох стандартних відхилень, і навпаки; якщо набір даних малий, дані можуть виходити за межі двох стандартних відхилень.

Якщо вхідні дані не розподілені за нормальним розподілом, то в цьому випадку метод середньоквадратичного відхилення не може бути використаний. Чудовим способом узагальнення розподілу гауссівських даних є міжквартильний розмах.

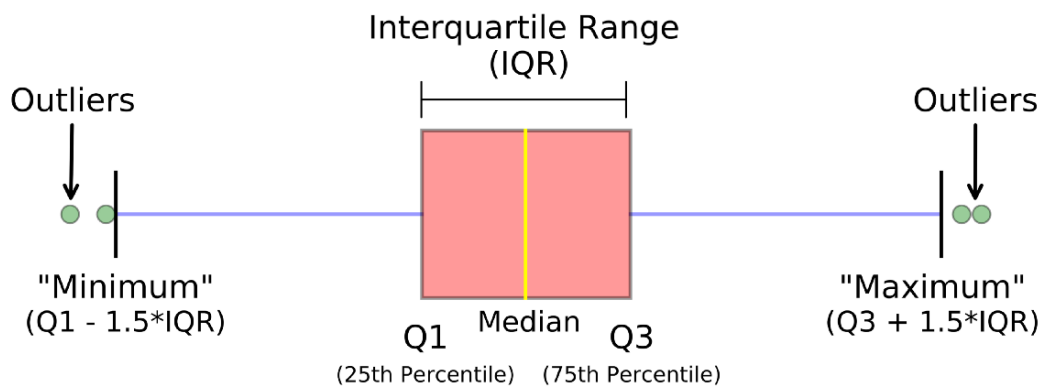


Рис. 2.7. Метод міжквартильного розмаху

Міжквартильний розмах розраховується як різниця між 75-м і 25-м процентилями даних. Статистичні методи виявлення викидів припускають, що викиди відбуваються в регіонах з низькою ймовірністю стохастичного розподілу, а отже, дані з'являються в регіонах з високою ймовірністю. Таким чином, метод міжквартильного розмаху може ідентифікувати викиди, визначаючи межі для вибірових значень, так званий параметр k IQR, тобто нижче 25-го перцентиля або вище 75-го перцентиля. Загальне значення коефіцієнта k становить 1,5. Будь-яке значення, що виходить за межі діапазону від $-1,5 \times \text{IQR}$ до $1,5 \times \text{IQR}$, вважається ненормальним (Рис. 2.7).

2.2.2. Вибір та аналіз моделей машинного навчання

Моделі часових рядів використовуються для прогнозування подій на основі перевірених історичних даних. Найпоширеніші типи включають ARIMA, згладжені та ковзні середні. Не всі моделі дають однакові результати для одного і того ж набору даних, тому дуже важливо визначити, яка з них працює найкраще, виходячи з індивідуальних часових рядів.

При прогнозуванні важливо розуміти свою мету. Щоб звузити специфіку вашої проблеми з прогнозним моделюванням, задайте питання про:

- Обсяг наявних даних - більше даних часто є більш корисними, пропонуючи більше можливостей для дослідницького аналізу даних, тестування та налаштування моделі, а також перевірки її точності.
- Необхідний часовий горизонт прогнозування - коротші часові горизонти часто легше прогнозувати - з більшою впевненістю - ніж

довші.

- Частота оновлення прогнозу - прогнози можуть потребувати частого оновлення з плином часу, а можуть бути зроблені один раз і залишатися статичними (оновлення прогнозів у міру надходження нової інформації часто призводить до більш точних прогнозів).
- Часова частота прогнозу - часто прогнози можуть бути зроблені на нижчих або вищих частотах, що дозволяє використовувати вибірку даних на нижчих і вищих частотах (це, в свою чергу, може дати переваги при моделюванні).

Прогнозування часових рядів - це використання моделі для передбачення майбутніх значень на основі попередніх спостережень [18]. Існує три аспекти прогнозного моделювання:

- Зразкові дані: дані, які ми збираємо, що описують нашу проблему з відомими взаємозв'язками між входами і виходами.
- Вивчення моделі: алгоритм, який ми використовуємо на основі даних вибірки для створення моделі, яку ми можемо використовувати знову і знову.
- Прогнозування: використання вивченої моделі на нових даних, для яких ми не знаємо результатів.

Серед факторів, які роблять прогнозування часових рядів складним, є наступні:

- Часова залежність часового ряду - основне припущення лінійної регресійної моделі про те, що спостереження є незалежними, у цьому випадку не спрацьовує. Через часову залежність даних часових рядів, прогнозування часових рядів не може покладатися на звичайні методи валідації. Щоб уникнути упереджених оцінок, навчальні набори даних повинні містити спостереження, які відбулися раніше, ніж ті, що використовуються у валідаційних наборах. Після того, як ми вибрали найкращу модель, ми можемо застосувати її до всієї навчальної вибірки та оцінити її ефективність на окремому тестовому наборі, який буде використано пізніше в часі.

- Сезонність у часовому ряді - Поряд з тенденцією до зростання або спадання, більшість часових рядів мають певну форму сезонних тенденцій, тобто варіацій, характерних для певного періоду часу.

Моделі часових рядів можуть перевершувати інші на певному наборі даних - одна модель, яка найкраще працює на одному типі даних, може не працювати так само для всіх інших.

Вибір моделей машинного навчання здійснювався опираючись на вже існуючі суміжні дослідження, а також опубліковану власну статтю з релевантним дослідженням, де проводився аналіз та працював із простими моделями, такими як лінійна та поліноміальна регресії, дерева рішень та модель випадкового лісу.

Проаналізувавши все вище сказане, мій вибір зупинився на дослідженні нейронних моделей CNN та LSTM, враховуючи їх нещодавній успіх у порівнянні з традиційними підходами [19].

Нейронні мережі – це основа алгоритмів глибинного навчання. Нейронні мережі ще називають штучні нейронні мережі. І назва полягає у тому що вони намагаються відтворити діяльність мозку людини, а саме мати структурно поєднанні нейрони, що передають між собою сигнали (Рис. 2.8).

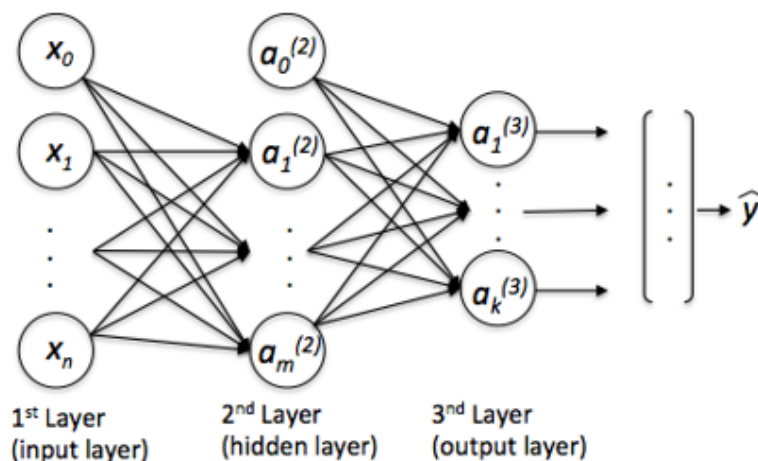


Рис. 2.8. Проста нейронна мережа

Проста модель нейронної мережі (персептрон) проводить процес навчання шляхом прямих і зворотних проходів, також відомих як пряме і зворотне поширення. Вибір гіперпараметрів у нейронній мережі має

вирішальне значення для управління її архітектурою, швидкістю навчання та регуляризацією, забезпечуючи адекватне навчання та запобігаючи надмірному пристосуванню.

Процес навчання в нейронній мережі (персептроні) включає наступні кроки:

- Архітектура моделі: включає кількість шарів, кількість нейронів у кожному шарі та функції активації, що використовуються в кожному нейроні. Для регресії вихідний шар зазвичай складається з одного нейрона з лінійною функцією активації (або без функції активації).
- Пряме поширення (Forward Propagation): процес поширення вхідних даних подається через нейронну мережу, шар за шаром, для генерації прогнозів. Активації нейронів у кожному шарі обчислюються за допомогою зважених сум виходів попереднього шару та відповідних функцій активації.
- Функція втрат: визначається для вимірювання різниці між прогнозованим виходом і фактичними цільовими значеннями. Зазвичай для задач регресії використовують середньоквадратичну похибку (MSE) або середню абсолютну похибку (MAE).
- Зворотне поширення (Backpropagation): оновлює ваги та зміщення нейронної мережі, щоб мінімізувати функцію втрат. Він обчислює градієнти втрат для параметрів мережі і використовує градієнтний спуск (або його різновиди) для оновлення ваг і зсувів у напрямку, який мінімізує втрати.
- Оновлення ваг та зсувів: отримуються шляхом множення градієнтів на швидкість навчання (розмір кроку) і віднімання їх від поточних ваг та зсувів.
- Повторення: процес прямого та зворотного поширення повторюються для декількох епох або до тих пір, поки модель не отримає задовільний розв'язок.

Гіперпараметри в нейронній мережі керують архітектурою моделі та процесом навчання. Вибір гіперпараметрів є важливим для успішного навчання

та продуктивності моделі. До найпоширеніших гіперпараметрів у нейромережевій моделі належать:

- Кількість шарів і нейронів: визначає архітектуру та складність мережі. Занадто мала кількість нейронів може не вловлювати складні закономірності, тоді як занадто велика кількість може призвести до надмірного пристосування. Налаштування гіперпараметрів може допомогти знайти оптимальний баланс.
- Функції активації: впливає на здатність мережі моделювати складні взаємозв'язки. Найпоширеніші варіанти включають ReLU (Rectified Linear Unit - випрямлена лінійна одиниця) та її різновиди. Лінійна функція активації зазвичай використовується для вихідного шару регресії.
- Швидкість навчання (Альфа): визначає розмір кроку для оновлення ваг під час градієнтного спуску. Висока швидкість навчання може призвести до перевищення оптимальних ваг, тоді як низька швидкість навчання може спричинити повільну збіжність.
- Розмір партії: визначає кількість зразків, що використовуються в кожному проході вперед і назад. Це може вплинути на швидкість навчання та використання пам'яті. Найпоширеніші варіанти: пакетний градієнтний спуск, міні-пакетний градієнтний спуск і стохастичний градієнтний спуск (SGD).
- Кількість епох: визначає кількість разів використання всього набору навчальних даних під час навчання. Занадто мала кількість епох може призвести до недостатнього навчання, а занадто велика - до надмірного.

2.2.3. Розбивання даних

Для оптимально навчання моделей прогнозування, набір даних розділяють на дві частини: навчальний набір, що включає перші 80% даних, і тестовий набір, що включає останні 20%. Кожен набір даних проходить процес сегментації на навчальний і тестовий набори, перш ніж його буде введено в модель машинного навчання. Навчальний набір даних відіграє вирішальну роль

у визначенні оптимальних параметрів моделі машинного навчання, що досягається за допомогою оптимізаторів, які зазвичай використовують такі методи, як градієнтний спуск. І навпаки, тестовий набір даних слугує засобом для неупередженої оцінки навченої моделі машинного навчання. Ця неупередженість зберігається завдяки тому, що модель не піддавалася впливу цього набору даних на етапі навчання. Використання стратегії створення валідаційних наборів даних ще більше покращує процес навчання. Ці валідаційні набори даних полегшують неупереджені оцінки на етапі навчання, сприяючи регуляризації, запобігаючи перенавчанню і потенційній деградації вихідної моделі, часто за допомогою механізмів ранньої зупинки.

Проте для прогнозування часових рядів, існують зовсім інші підходи. Далі розглянемо один із них, а саме метод вікон (Рис. 2.9).

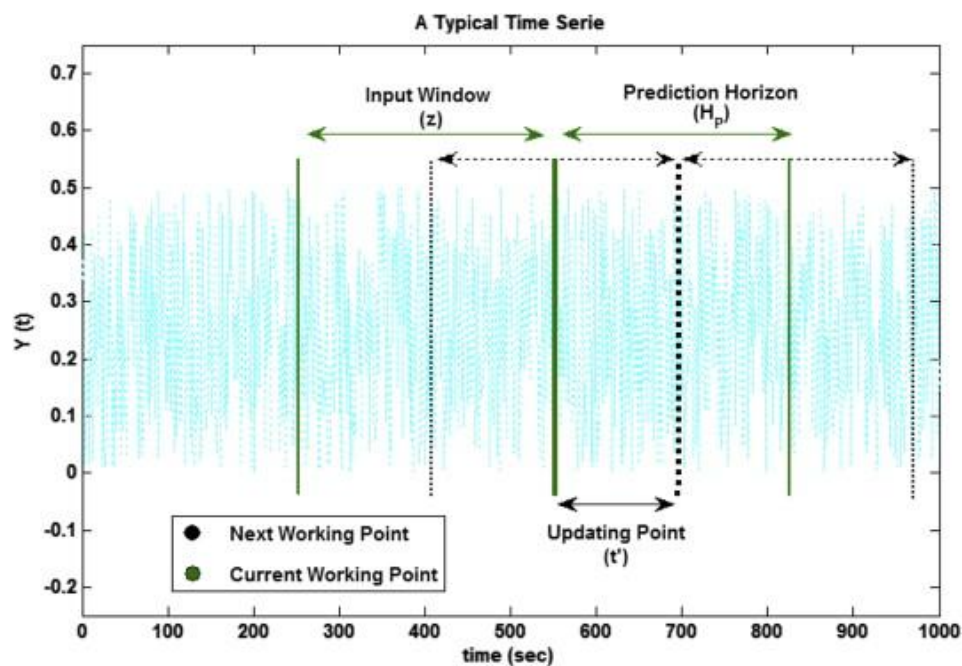


Рис. 2.9. Схематичне представлення роботи методу вікон для часових рядів

Метод вікон є поширеним підходом [18], який використовується в прогнозуванні часових рядів для навчання моделей машинного навчання. Він передбачає поділ даних часового ряду на послідовні сегменти, що не перекриваються, або "вікна", і використання цих сегментів як вхідних даних для навчання моделі. Цей метод особливо корисний, коли ви хочете виявити часові закономірності та залежності в даних.

Ось покроковий опис методу вікон:

- Підготовка даних: почніть з даних часового ряду, який зазвичай являє собою послідовність спостережень, зібраних протягом певного часу. Визначте розмір вікна, тобто кількість часових кроків, які потрібно включити в кожен сегмент. Це дуже важливий параметр, і його слід обирати, виходячи з характеристик ваших даних і закономірностей, які ви хочете, щоб модель відобразила.
- Створення вікон: створіть вікна даних, що перекриваються або не перекриваються, пересуваючи вікно по часовому ряду. Наприклад, якщо ваш часовий ряд має вигляд [1, 2, 3, 4, 5, 6, 7, 8, 9] і ви вибрали розмір вікна 3, ви створите такі вікна: [1, 2, 3], [4, 5, 6], [7, 8, 9].
- Пари вводу-виводу: для кожного вікна визначте частину даних як вхідні характеристики, а наступний часовий крок - як вихідний або цільовий. Наприклад, використовуючи розмір вікна 3, парами вхідних даних для наведеного вище прикладу будуть ([1, 2], 3), ([4, 5], 6), ([7, 8], 9).
- Навчання моделі: використовуйте ці пари вхідних-вихідних даних для навчання вашої моделі прогнозування часових рядів. Модель вчиться передбачати наступний крок на основі шаблонів, присутніх у вікні введення.
- Прогнозування: після навчання ви можете використовувати модель для прогнозування нових даних. Щоб передбачити наступний часовий крок, ви використовуєте ковзне вікно над невидимою частиною часового ряду, подібно до того, як ви створювали навчальні вікна.

Метод вікон дозволяє моделі вивчати закономірності та залежності в даних часового ряду, що робить його цінною технікою для задач прогнозування часових рядів. Важливо експериментувати з різними розмірами вікон, щоб знайти оптимальний для вашого конкретного набору даних і цілей прогнозування.

2.2.4. Масштабування набору даних

Оскільки набори даних мають широкий діапазон значень, набір даних трансформується, тобто нормалізується наступним чином, щоб допомогти структурувати процес навчання та побудувати покращену модель. Тому важливо їх нормалізувати. Масштабування ознак - це метод, який приводить значення даних до коротшого діапазону.

Наприклад, одна характеристика може мати діапазон вимірювання від 0,0001 до 0,2, а інша - від -100 до 100. Наприклад, вік клієнта може бути від 16 до 40 років, але більшість клієнтів віком від 18 до 25 років, тому математичне очікування зміщується до центру розподілу. Ця характерна відмінність може спричинити значні помилки в багатьох моделях (наприклад, для регресії, нейронних мереж). Тому необхідно інтегрувати всі функції в одну форму.

Існує певна плутанина щодо термінів "стандартизація" та "нормалізація" [1]. Стандартизація і нормалізація часто розглядаються як різні речі, а іноді як частина нормалізації. Тому важливо розуміти загальну природу та призначення цих методів.

Стандартизація даних - це процес надання вектору кожного атрибуту такого вигляду, щоб його математичне сподівання стало нульовим, а дисперсія - одиничною.

Розглянемо, як можна масштабувати дані для їхньої стандартизації:

- Масштабування з використанням мінімальних і максимальних значень (Формула 2.1).

$$X_i^{scaled} = \frac{X_i - \min(X)}{\max(X) - \min(X)} \quad (2.1)$$

Цей підхід масштабує кожну змінну таким чином, щоб вона була у межах певного діапазону навчального набору даних, наприклад, від 0 до 1.

- Масштабування змінних за максимальною абсолютною величиною (Формула 2.2).

$$X_i^{scaled} = \frac{X_i}{\max(|X|)} \quad (2.2)$$

Дана функція масштабує кожну змінну, щоб максимальне абсолютне значення кожного атрибута у навчальному наборі дорівнювало 1. Так як алгоритм не переміщує і не центрує дані, вінповідно не порушує розрідженості.

- Масштабування змінної, до середньоквадратичного відхилення (Формула 2.3).

$$X_i^{scaled} = \frac{X_i - mean(X)}{std(X)} \quad (2.3)$$

Нормалізація даних — це процес масштабування вектора кожної ознаки. Щоб зробити всі вектори однакового масштабу, потрібна нормалізація. Існує кілька типів норм і нормалізацій:

- Максимальна норма (Формула 2.4)

$$X_i^{norm} = \frac{X_i}{max(X)} \quad (2.4)$$

Щоб всі значення лежали в межах діапазону, потрібно знайти максимально можливе значення і розділити на нього всі інші значення. Тому максимальним значенням буде одиниця, а всі інші потрапляють в діапазон від 0 до 1. При умові, якщо немає негативних значень.

- L1-norm (Формула 2.5)

$$X_i^{norm} = \frac{X_i}{\sum_{i=1}^m X_i} \quad (2.5)$$

Кожне значення потрібно поділити на суму значень заданого розподілу.

- L2-norm (Формула 2.6)

$$X_i^{norm} = \frac{X_i}{\sqrt{\sum_{i=1}^m X_i^2}} \quad (2.6)$$

Даний підхід також називають евклідовою нормою.

Масштабування атрибутів є кінцевим етапом попередньої обробки даних у машинному навчанні. Цей метод приводить незалежні змінні набору даних у певний загальний діапазон значень. Іншими словами, масштабування атрибутів обмежує діапазон змінних, щоб модель могла порівнювати їх з усіма іншими, незалежно від масштабу значення.

2.2.5. CNN

У сфері комп'ютерного зору та автоматизованого розпізнавання мови (ASR) однією з найвідоміших моделей глибокого навчання є згортова нейронна мережа (CNN). Хоча CNN широко відомі своєю ефективністю в розпізнаванні зображень і відео, рекомендаційних системах, аналізі медичних зображень і обробці природної мови, вони також довели свою цінність в обробці даних часових рядів. Фундаментальна архітектура CNN для роботи з часовим рядом представлена на рисунку 2.10.

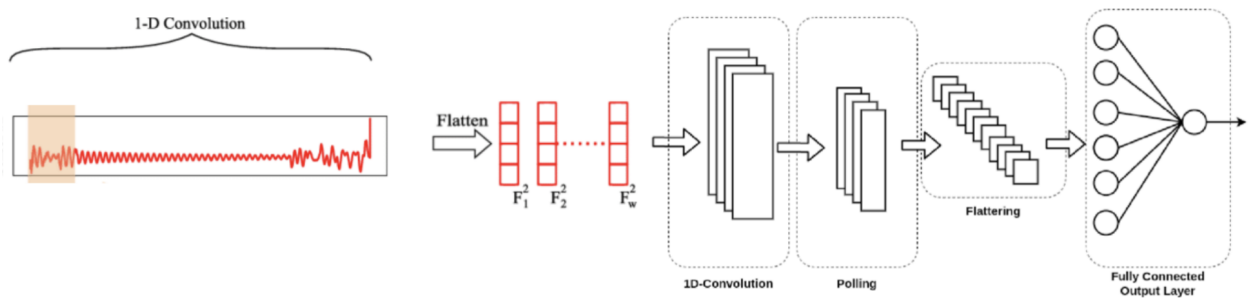


Рис. 2.10. Схема моделі CNN для роботи з часовими рядами

Однією з важливих характеристик CNN є їхня здатність обробляти багатоканальні вхідні дані, що робить їх особливо придатними для обробки різних часових рядів даних з кількома входами і виходами [19]. Проте, залишається недостатньо досліджень щодо ефективності CNN у моделюванні та прогнозуванні значень множинних часових рядів даних в рамках моделей глибокого навчання.

Ключовою перевагою конволюційної нейронної мережі є те, що вона чудово розпізнає локальні закономірності в даних. У контексті часових рядів це означає здатність мережі ідентифікувати часові закономірності та залежності в послідовності. Наприклад, в аналізі фінансових даних ШНМ можуть фіксувати локальні тенденції та аномалії в цінах на акції.

Конволюційної нейронної мережі використовують розподіл ваги, що означає, що один і той самий набір параметрів (фільтрів), що навчаються, застосовується до різних областей вхідних даних. Ця властивість дозволяє CNN ефективно навчатися і розпізнавати схожі патерни на різних ділянках часового

ряду. Наприклад, у розпізнаванні мови, де фонемі виникають на різних часових кроках, розподіл ваги допомагає виявити повторювані акустичні ознаки.

Архітектури CNN зазвичай складаються з декількох згорткових шарів і шарів об'єднання. Ці шари працюють разом для поступового вилучення абстрактних та ієрархічних ознак з вхідного часового ряду. Ранні шари фіксують низькорівневі ознаки, такі як краї або короткострокові патерни, тоді як більш глибокі шари фіксують більш складні, високорівневі часові ознаки. Щоб пом'якшити проблему, пов'язану зі збільшенням розмірності ознак, і зменшити обчислювальне навантаження на процес навчання, стратегічно включено шар дропауту. Його основна мета - зменшити кількість вилучених ознак перед тим, як вони будуть використані в кінцевому результаті.

Плюси та мінуси CNN для аналізу часових рядів.

Плюси:

- Виділення ознак: CNN мають досвід автоматизованого вилучення ознак, що зменшує потребу в ручній розробці ознак для аналізу часових рядів.
- Ефективні локальні закономірності: вони чудово виявляють локальні закономірності та залежності, що робить їх придатними для завдань, де короткострокові часові зв'язки мають вирішальне значення.
- Розподіл ваги: розподіл ваги значно зменшує кількість параметрів, що робить ШНМ обчислювально ефективними, особливо для великих наборів даних часових рядів.
- Перенесення навчання: попередньо навчені моделі ШНМ з інших областей, наприклад, розпізнавання зображень, можуть бути точно налаштовані для роботи з часовими рядами, використовуючи знання, отримані з різних джерел даних.

Недоліки:

- Послідовна інформація: ШНМ можуть не вловлювати довготривалі послідовні залежності так само ефективно, як рекурентні нейронні мережі (RNN) або трансформатори. Таким чином, вони можуть бути не

найкращим вибором для задач, що вимагають моделювання великих часових лагів.

- Фіксований розмір вхідних даних: ШНМ часто вимагають фіксованого розміру вхідних даних, що може бути обмежуючим фактором для часових рядів різної довжини.
- Складність: Розуміння та налаштування архітектури ШНМ для задач з часовими рядами може бути складним завданням, особливо для новачків у глибокому навчанні.

Отже, згорткові нейронні мережі стали універсальним інструментом для аналізу часових рядів, здатним ефективно виявляти локальні закономірності та залежності. Однак їхня придатність залежить від конкретних характеристик даних часових рядів і поставленого завдання. Дослідники і практики повинні ретельно зважити всі "за" і "проти" при виборі ШНМ для роботи з часовими рядами.

2.2.6. LSTM

Окрім CNN, модель довготривалої короткочасної пам'яті, також відома як LSTM (Рис. 2.11), є ще однією відомою моделлю DNN. LSTM - це вид рекурентної нейронної мережі (RNN). Алгоритм LSTM ідеально підходить для класифікації, сортування та прогнозування на основі одного набору даних часового ряду [19].

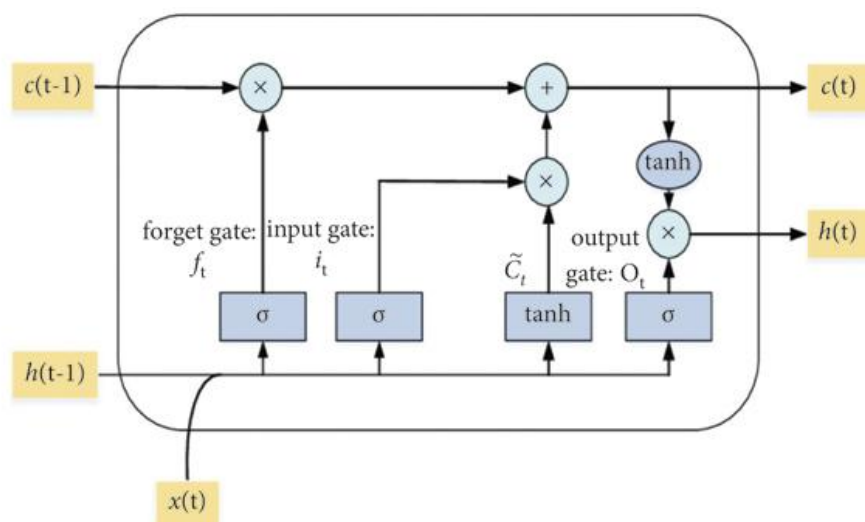


Рис. 2.11. Схема моделі LSTM

LSTM відрізняються своєю здатністю ефективно обробляти дані часових рядів завдяки наступним механізмам:

- **Послідовне керування пам'яттю:** LSTM мають унікальну архітектуру, яка дозволяє їм зберігати та маніпулювати інформацією в довгих послідовностях. Це досягається за допомогою набору спеціалізованих вентилів у комірці LSTM, які включають вхідний вентиль, вентиль забування та вихідний вентиль. Ці вентиля керують потоком інформації, дозволяючи мережі запам'ятовувати релевантні минулі події та забувати нерелевантну інформацію, що є важливою особливістю для фіксації часових залежностей у даних часових рядів.
- **Обробка послідовностей змінної довжини:** на відміну від багатьох інших нейромережових архітектур, LSTM можуть легко обробляти послідовності різної довжини. Ця гнучкість є безцінною для задач, де точки даних збираються через нерегулярні інтервали часу або при роботі з даними часових рядів, які містять пропуски або відсутні значення.
- **Механізми вентилявання:** основна перевага LSTM полягає в їхніх механізмах вентилявання. Ці механізми дозволяють мережі вирішувати, яку інформацію слід зберігати, яку відкидати, а яку оновлювати на кожному часовому кроці. Така адаптивність дозволяє LSTM навчатися і фіксувати складні закономірності в даних, враховуючи як короткострокові, так і довгострокові залежності.
- **Навчання з урахуванням станів:** LSTM можна навчати з урахуванням стану, коли внутрішній стан в кінці однієї послідовності можна використовувати як початковий стан для наступної послідовності. Ця характеристика особливо цінна для задач прогнозування часових рядів, де прогнози залежать від історичного контексту даних.
- **Паралельність:** хоча навчання LSTM може бути обчислювально інтенсивним, вони пропонують певний паралелізм під час навчання завдяки своїй послідовній природі. Пакетна обробка та сучасне обладнання, таке як графічні процесори, можуть бути використані для

скорочення часу навчання.

- Стекові LSTM: для більш складних задач обробки часових рядів кілька шарів LSTM можна накладати один на одного, створюючи глибокі LSTM-мережі. Ця ієрархічна архітектура дозволяє моделі охоплювати все більш абстрактні та складні часові характеристики.

Обчислювальний процес, який відбувається в структурі LSTM для обчислення прогнозних даних часових рядів, починається з обчислення вихідного значення з попереднього часу (Формула 2.7), а вхідне значення з поточного часу стає входом для воріт забування, і результати обробки воріт забування отримуються шляхом обчислення за наступною формулою:

$$f_t = \sigma(W_f * [h_{t-1}, x_t] + b_f), \quad (2.7)$$

де f_t - діапазон значень (0, 1), W_f - вага вентиля забування, b_f - значення зсуву, застосоване до вентиля забування, x_t - вхідне значення для поточного часу, h_{t-1} - вихідне значення попереднього часу обробки.

Крім того, вихідне значення попереднього часу та вхідне значення поточного часу також подаються на вхідні ворота (Формула 2.8), а вихідне значення та стан комірки-кандидата на вхідних воротах отримуються після обчислення формулою 2.9.

$$i_t = \sigma(W_i * [h_{t-1}, x_t] + b_i), \quad (2.8)$$

$$\hat{C}_t = \tanh * (W_c * [h_{t-1}, x_t] + b_c), \quad (2.9)$$

де діапазон його значень (0, 1), W_i - вага входу вентиля, b_i - значення зсуву входу вентиля, W_c - вага кандидата на вхід вентиля, b_c - значення зсуву кандидата на вхід вентиля.

Наступним етапом LSTM-моделі є процес коригування значень комірок або параметрів моделі в цей час, який здійснюється наступним чином (Формула 2.10):

$$C_t = f_t * C_{t-1} + i_t * \hat{C}_t, \quad (2.10)$$

де діапазон значень для C_t - (0, 1). Потім, в момент часу обробки t , вихідне

значення h_{t-1} і вхідне значення x_t стають входом для вихідного вентиля, а вихід з виходу вентиля обчислюється за наступною формулою 2.11:

$$o_t = \sigma * (W_0 * [h_{t-1}, x_t] + b_0), \quad (2.11)$$

де діапазон значень o_t становить $(0,1)$, W_0 - вага виходу вентиля, а b_0 - значення зсуву виходу вентиля.

Нарешті, остаточне вихідне значення LSTM генерується вихідним вентилям і є результатом розрахунку за наступною формулою 2.12:

$$h_t = o_t * \tanh(C_t), \quad (2.12)$$

У цьому випадку $\tanh$ - це функція активації, яка може бути адаптована до потреб і характеристик задачі, яку потрібно вирішити.

Отже, моделі з довготривалою короткочасною пам'яттю (LSTM) є чудовим вибором для аналізу часових рядів, оскільки вони здатні вловлювати складні часові залежності і враховувати послідовності різної довжини. Однак складність LSTM-моделей і необхідність ретельного налаштування гіперпараметрів слід враховувати при використанні їх для конкретних застосувань часових рядів.

2.3. Інтерфейс програми

Останнім компонентом системи є користувацький інтерфейс, який надає користувачам можливість самостійно оцінювати стан обладнання та отримувати персоналізовані рекомендації щодо оптимальної роботи обладнання у зручній для користувача формі.

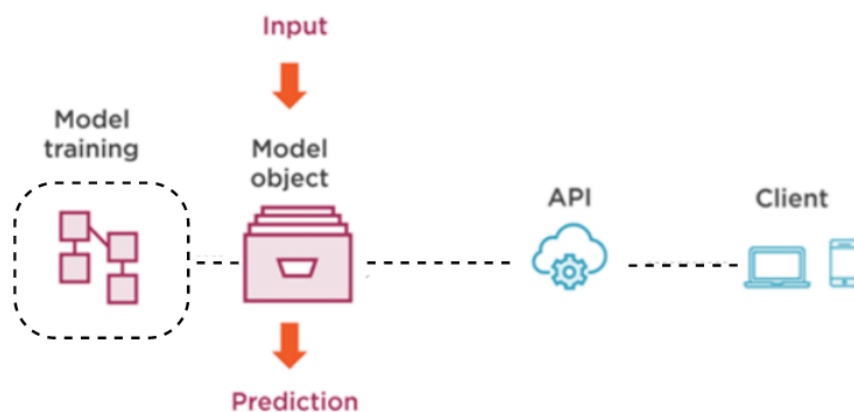


Рис. 2.12. Схема клієнт-сервісної архітектури системи

Розроблений за архітектурою клієнт-сервер (див. Рис. 2.12), цей модуль містить інтуїтивно зрозумілі та доступні інструменти. Вхідні параметри включають ідентифікацію конкретної одиниці важкої техніки та деталі щодо характеру і складності очікуваних завдань. У відповідь користувацький інтерфейс видає індивідуальні рекомендації щодо використання техніки, згенеровані завдяки застосуванню навчених моделей машинного навчання. Такий дизайн, орієнтований на користувача, має на меті дати можливість людям приймати обґрунтовані рішення щодо використання обладнання, підвищуючи операційну ефективність і подовжуючи загальний термін служби техніки.

2.4. Висновки до розділу

Прогнозування часових рядів - це процес аналізу даних часових рядів з використанням статистики та моделювання для складання прогнозів та прийняття стратегічних рішень. Важливою відмінністю прогнозування є те, що на момент виконання роботи майбутній результат повністю недоступний і може бути оцінений лише за допомогою ретельного аналізу та науково обґрунтованих попередніх даних. Це не завжди точний прогноз, і ймовірність прогнозів може сильно відрізнятися - особливо, коли йдеться про змінні, що часто змінюються в даних часових рядів, а також про фактори, які підлягають контролю. Однак прогнозування дає уявлення про те, які результати є більш ймовірними - або менш ймовірними - порівняно з іншими потенційними результатами.

Прогнозування часових рядів не є безпомилковим і не є доречним або корисним для всіх ситуацій. Оскільки не існує чіткого набору правил щодо того, коли слід чи не слід використовувати прогнозування, аналітики та команди, що працюють з даними, повинні знати обмеження аналізу і те, що можуть підтримувати їхні моделі. Не кожна модель підійде до кожного набору даних або відповідь на кожне питання. Команди даних повинні використовувати прогнозування часових рядів, коли вони розуміють бізнес-питання і мають відповідні дані та можливості прогнозування, щоб відповісти на це питання. Хороше прогнозування працює з чистими даними з позначкою часу і може виявити справжні тенденції та закономірності в історичних даних. Аналітики

можуть відрізнити випадкові коливання від випадкових відхилень і відокремити справжні тенденції від сезонних коливань. Аналіз часових рядів показує, як дані змінюються з часом, а хороше прогнозування може визначити напрямок, в якому ці дані змінюються.

Випадкові речі ніколи не будуть точно спрогнозовані, незалежно від того, скільки даних збирається або наскільки послідовно. Наприклад: можна щотижня спостерігати за даними про кожного переможця лотереї, але передбачити, хто виграє наступного разу ніяк не можна. Зрештою, саме ваші дані та аналіз часових рядів даних визначають, коли вам слід використовувати прогнозування, оскільки прогнози дуже різняться залежно від різних факторів. Використовуйте своє судження і знайте свої дані. Пам'ятайте про цей список міркувань, щоб завжди мати уявлення про те, наскільки успішним буде прогнозування.

3. РОЗДІЛ АПРОБАЦІЙ ТА РЕЗУЛЬТАТІВ

3.1. Засоби реалізації

При програмній реалізації системи, призначеної для генерування персоналізованих рекомендацій для важкої техніки, з урахуванням важливих характеристик, було зроблено стратегічний вибір технологій розробки, в першу чергу, з використанням Python та JavaScript.

Python, відомий своєю практичністю в аналізі даних та використанні інструментів машинного навчання, слугує основою реалізації. Серед ключових бібліотек, які використовуються в цьому проекті, - Pandas, відома своїми зручними функціями для аналізу та маніпулювання даними; NumPy, що допомагає працювати з даними у вигляді масивів; SciPy, що використовується для математичного та числового аналізу; Scikit-Learn, широко використовувана бібліотека для науки про дані та машинного навчання; а також TensorFlow і Keras.

Pandas спрощує трудомісткі та повторювані завдання, що є невід'ємною частиною управління даними, охоплюючи такі функції, як очищення даних, заповнення прогалів, злиття, нормалізація, візуалізація, а також завантаження та зберігання даних.

Тим часом, включення бібліотеки Scikit-Learn полегшує виконання багатьох операцій і пропонує широкий спектр алгоритмів машинного навчання. Хоча вона кваліфіковано підтримує такі завдання, як попередня обробка даних, зменшення розмірності, вибір моделі, регресія та класифікація, слід зазначити, що вона не забезпечує вичерпної підтримки нейронних мереж. Для цього в систему інтегровані TensorFlow і Keras, які використовують можливості нейронних мереж для більш складних завдань машинного навчання.

Побудова нейромережевої моделі полегшується завдяки використанню пакетів TensorFlow та Keras. Вибір саме цих пакетів для розробки обумовлений декількома вагомими причинами:

- Вичерпна документація: TensorFlow і Keras пропонують детальну документацію, що надає користувачам вичерпні вказівки щодо ефективного використання їхніх методів.
- Крос-платформне розгортання: пакети дозволяють безперешкодно розгортати моделі на різних мовах і платформах, підвищуючи універсальність розробки додатків.
- Широка підтримка моделей: TensorFlow і Keras можуть похвалитися сумісністю з різноманітними моделями машинного навчання, що сприяє гнучкості у виборі моделей.
- Зручне налаштування навчання: пакети пропонують зручні інструменти для налаштування навчання моделей, що спрощує процес кастомізації для розробників.

При створенні клієнт-серверної структури для розгортання бекенду було обрано фреймворк Flask, інструмент на основі Python. Водночас для створення користувацького інтерфейсу було обрано фреймворк Angular, технологію на основі JavaScript. Ця стратегічна комбінація забезпечує надійну та згуртовану систему, в якій Flask виконує внутрішні операції, а Angular забезпечує інтуїтивно зрозумілий користувацький інтерфейс.

3.2. Вимоги до технічного та програмного забезпечення

До функціональних вимог системи розробки відноситься, перш за все, загальний вигляд програми: це сторінка користувача, на якій перераховані екземпляри важкої техніки та типи складності роб, які користувач повинен обрати, щоб отримати прогнозовану оцінку щодо експлуатації обраної техніки. Прогнозування має повернути текстове повідомлення, обрана користувачем техніка завершить роботу успішно без поломи, або ж під час виконання поставлених цілей отримає поломку, яка потенційно вплине на успіх та терміни виконання завдання. Якщо користувач не обрав вид техніки та складність роботи, то система повідомить йому про це. Після натискання кнопки «Отримати оцінку» програма поверне текстове повідомлення про

виконання даного завдання, а також значення прогнозування даних у вигляді графіків, для кращого візуального сприйняття даних користувачем

До нефункціональних вимог програми в основному відноситься надійність, тобто система повинна реально оцінювати надані їй параметри. Система повинна бути простою у використанні – це також одна з вимог, зокрема, реалізований графічний інтерфейс користувача, що дає змогу будь-кому використовувати систему.

Вимогами до програмного забезпечення:

- Мова програмування Python (Flask) , JavaScript (Angular)
- Windows, Linux OS, Mac OS, чи будь яка інша операційна система
- База даних PostgreSQL

Вимогами до технічного обладнання:

- Ноутбук або ПК
- ОПЗ 6 Гб та більше
- Стабільне з'єднання з інтернетом

3.3. Опис програмної реалізації

Одним із основних етапом робробки системи це є збір даних та їх обробка та структуризація. До даного процесу слід ретельно підійти, так як від якості даних, залежить складність їх трансформації для моделі навчання, а також оцінки тренування моделей машинного навчання. Для дослідження прогнозування показів датчиків у вигляді часових рядів, було зібрано датасет, який налічує понад 50 різних датчиків. Загальний набір даних до обробки складав понад 220 тисяч записів часового ряду.

Таблиця 3.1. Опис набору даних

Кількість датчиків		Кількість записів	
До обробки	Після обробки	До обробки	Після обробки
51	46	220320	114050

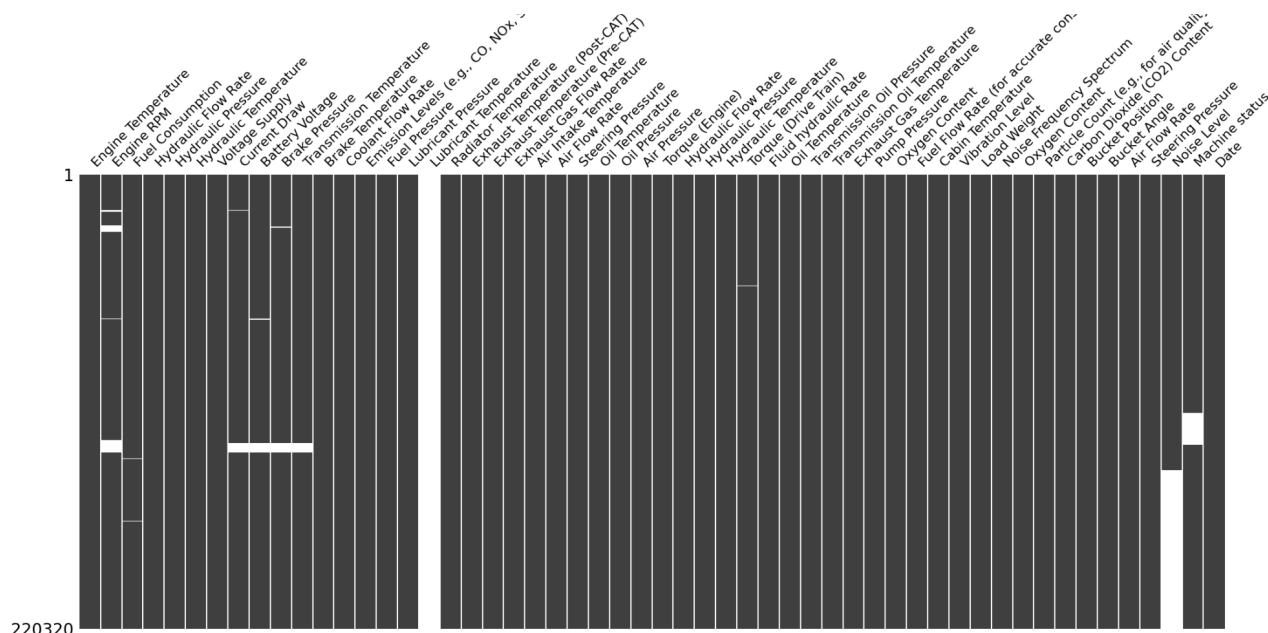


Рис. 3.1. Вигляд даних до процесу обробки

В таблиці 3.1 представлено кількість датчиків та кількість записів по кожному із датчиків до та після обробки. На рисунку 3.1 зображено щільність заповнення значень кожного датчика. З результатів видно, що в процесі препроцесингу було відкинуто близько половини усіх значень по кожному датчику, проте в нас залишився масивний набір очищених даних, який буде використовуватися в дослідженні.

Перш ніж приступати до побудови та навчання моделей машинного навчання, необхідно провести попередню обробку даних. Основні етапи попередньої обробки включають виявлення та обробку відсутніх значень, виключення аномальних даних та масштабування набору даних.

Після ретельного аналізу даних на наявність пропусків було прийнято рішення видалити їх з набору даних, що було обґрунтовано значним розміром вибірки.

Для виявлення аномальних даних було використано два загальновизнані методи: метод міжквартильного розмаху та метод стандартного відхилення. Слід зазначити, що застосування методу стандартного відхилення залежить від того, чи відповідають дані нормальному розподілу. Щоб переконатися в цьому, були використані відомі тести (такі як тест Колмогорова-Смірнова та тест Ліліфорса). Ці тести оцінюють розподіл даних і надають булеве значення, яке

вказує, чи відповідають дані нормальному розподілу. На основі результатів цих тестів було застосовано метод видалення аномальних даних. Варто зазначити, що оцінка аномальних даних проводилася незалежно для кожного набору згрупованих даних з однаковими параметрами.

	Engine Temperature	Engine RPM	Fuel Consumption	Hydraulic Flow Rate	Hydraulic Pressure	Hydraulic Temperature	Voltage Supply	Current Draw	Battery Voltage	Brake Pressure	...
count	114050.000000	114050.000000	114050.000000	114050.000000	114050.000000	114050.000000	114050.000000	114050.000000	114050.000000	114050.000000	...
mean	29.605708	29.579045	29.595513	29.591973	29.609161	29.472324	29.581396	29.604375	29.605315	29.609711	...
std	0.557194	1.374248	0.956042	1.059995	0.325597	2.861920	1.322660	0.624026	0.577708	0.270962	...
min	27.131703	24.063833	25.760825	25.254901	28.140281	19.554554	24.375773	26.197035	27.668850	28.468168	...
25%	29.463705	28.609223	28.934439	28.887657	29.404895	27.414349	28.476817	29.319498	29.121468	29.547263	...
50%	29.593921	29.678726	29.628669	29.701392	29.615664	29.444432	29.045123	29.570880	29.633341	29.604067	...
75%	30.091093	30.507593	30.273308	30.369819	29.821026	31.369921	30.872919	30.047196	30.020712	29.675042	...
max	30.671134	34.411278	32.008882	32.287918	30.966744	38.866530	33.883404	31.687816	31.750007	30.555364	...

8 rows x 45 columns

Рис. 3.2. Статистичні деталі набору даних

Метод ".describe" з бібліотеки Pandas (див. Рис. 3.2) виявився корисним для отримання важливих статистичних даних. Цей метод полегшує вивчення фундаментальних статистичних деталей, включаючи процентилі, середнє значення, стандартне відхилення, а також максимальні та мінімальні значення.

Date	Hydraulic Temperature(x-5)	Hydraulic Temperature(x-4)	Hydraulic Temperature(x-3)	Hydraulic Temperature(x-2)	Hydraulic Temperature(x-1)	Hydraulic Temperature(x)	Hydraulic Temperature(x+60)
2023-06-14 19:00:00	26.363787	28.645428	26.447322	25.779618	27.606078	25.427968	24.938284
2023-06-14 19:01:00	27.466801	28.729869	25.364832	25.637739	26.829366	25.397170	25.362597
2023-06-14 19:02:00	27.458789	27.937468	25.681190	25.422530	25.931836	25.226620	24.870360
2023-06-14 19:03:00	27.381493	27.192164	26.367888	25.594479	26.159407	25.621499	24.098316
2023-06-14 19:04:00	27.952406	27.773798	27.072404	25.515495	25.822831	25.016716	23.649557
...	...	...	...	...	...	...	...
2023-06-15 22:55:00	31.983332	29.639826	27.088380	29.415618	28.318291	25.803302	23.234612
2023-06-15 22:56:00	31.900438	30.655772	27.088380	29.504728	27.673676	25.602945	22.491994
2023-06-15 22:57:00	32.586074	30.788977	26.648371	28.869994	27.057578	24.836355	23.005136
2023-06-15 22:58:00	32.365827	31.639025	26.349225	29.434210	26.846788	25.436376	23.815811
2023-06-15 22:59:00	31.818559	31.236737	25.243284	28.562145	26.421293	26.243362	23.661308

Рис. 3.3. Представлення роботи методу вікон для часових рядів на основі одного із датчиків важкої техніки (крок вікна 60, що еквівалентно 60 хвилинам в даному числовому ряді)

Наданий набір даних являє собою часовий ряд показів температури гідравлічної системи, отриманих з датчика важкої техніки. Для ефективного аналізу та обробки цих часових даних застосовано метод вікон, де кожне

спостереження фіксує гідравлічну температуру в різні моменти часу. Набір даних організовано в хронологічному порядку, де кожен рядок відповідає певній часовій позначці.

Для реалізації методу вікон обрано розмір вікна 60, що означає проміжок часу в 60 хвилин. Цей метод передбачає створення послідовних сегментів або "вікон" даних часового ряду, кожен з яких складається з фіксованої кількості послідовних спостережень. У цьому випадку кожне вікно інкапсулює показники через регулярні інтервали, що охоплюють одну годину.

Наприклад, перше вікно починається з початкової часової позначки і включає точки даних до 60 хвилин потому. Наступні вікна продовжують просуватися з кроком 60 хвилин, фіксуючи зміну показників гідравлічної температури в часі. Такий підхід полегшує аналіз тенденцій і закономірностей у кожному часовому сегменті, пропонуючи розуміння поведінки обладнання.

Представлена рисунку 3.3 ілюструє результат застосування цього методу до даних часових рядів. Столпчики "Гідравлічна температура (x-5)" - "Гідравлічна температура (x+60)" представляють показники в кожному вікні, забезпечуючи структурований і організований вигляд варіацій гідравлічної температури протягом послідовних 60-хвилинних інтервалів. Цей метод допомагає виявити часові закономірності і полегшує вилучення значущої інформації з даних датчиків важкого обладнання.

3.3.1. Аналіз результатів

Для оцінки точності прогнозування зазвичай використовують одну або кілька метрик якості, які часто зосереджуються на розбіжності між розрахованою відповіддю та очікуваною, по суті, кількісно оцінюючи обчислювальні помилки. Альтернативно, інший метод оцінки якості прогнозування передбачає обчислення коефіцієнта детермінації.

Середня абсолютна похибка (MAE) визначається шляхом підсумовування абсолютних різниць між фактичними і прогнозованими значеннями для кожного запису в наборі даних, а потім ділиться на загальну кількість значень у масиві (Формула 3.1).

$$MAE = \frac{\sum_{i=1}^n |y_i - \bar{y}_i|}{n} \quad (3.1)$$

Середньоквадратична похибка (MSE) оцінюється як середньоквадратична різниця між фактичними і прогнозованими значеннями за формолою 3.2.

$$MSE = \frac{\sum_{i=1}^n (y_i - \bar{y}_i)^2}{n} \quad (3.2)$$

Середньоквадратична похибка (RMSE) обчислюється як квадратний корінь із середньоквадратичної похибки (Формула 3.3), що відображає концентрацію даних відносно лінії найкращого узгодження. Цей метод дає уявлення про те, наскільки точно точки даних узгоджуються з оптимальною моделлю прогнозування.

$$RMSE = \sqrt{\frac{\sum_{i=1}^n (y_i - \bar{y}_i)^2}{n}} \quad (3.3)$$

Для кращого візуального сприйняття, було прийнято рішення зобразити значення похибок у вигляді графіків, числові значення похибок по кожному датчику відповідної моделі представленні в додатку А (див. Таблиця А.1).

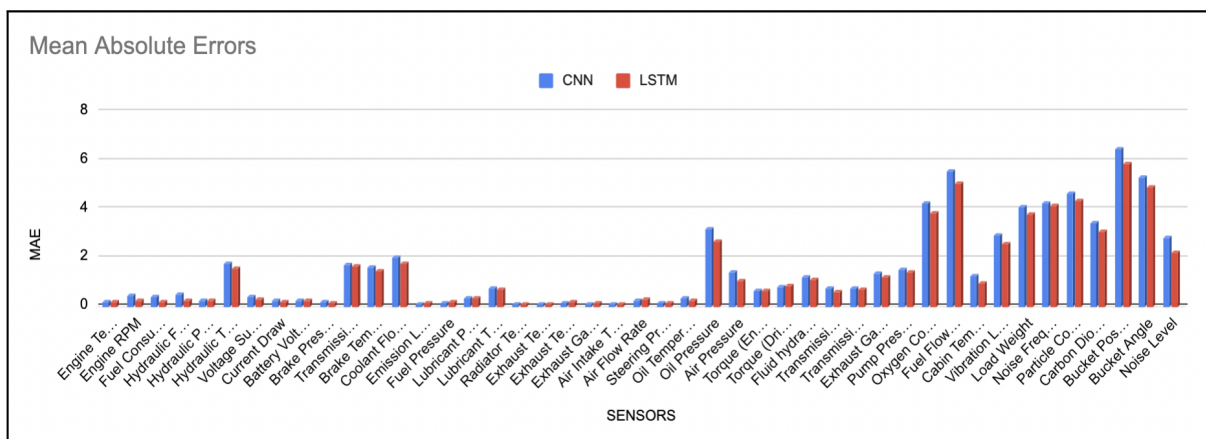


Рис. 3.4. Значення похибки MAE

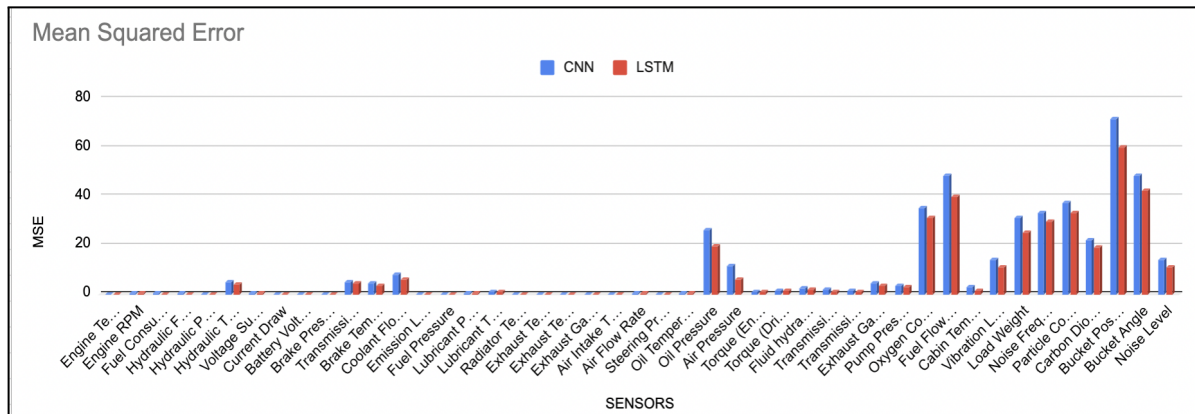


Рис. 3.5. Значення похибки MSE

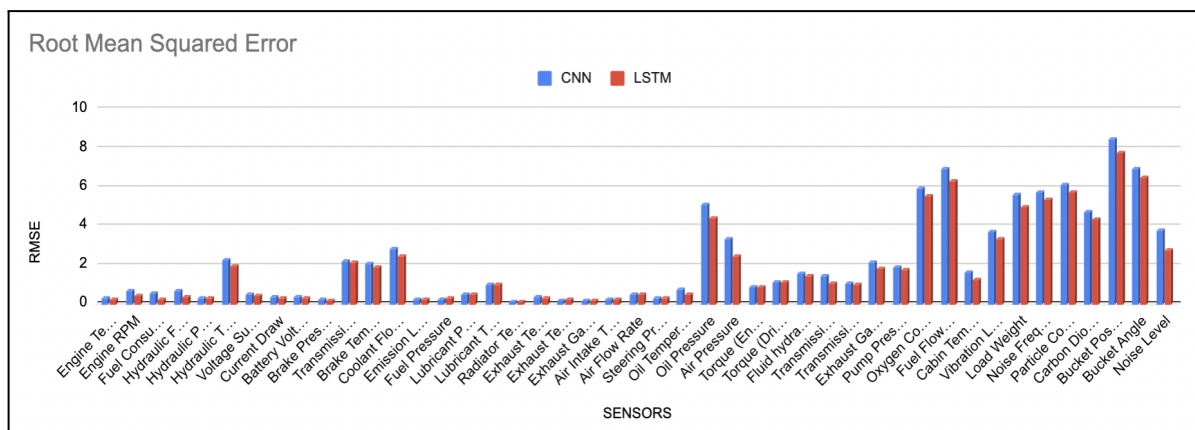


Рис. 3.6. Значення похибки RMSE

Аналіз результатів показує, що модельні структури, розроблені для цього дослідження, принесли частковий успіх. Ретельно розроблені архітектури нейронних мереж продемонстрували хорошу продуктивність приблизно на половині датчиків, що свідчить про значне досягнення. Цей результат підкреслює ефективність запропонованої архітектури моделі у виявленні та прогнозуванні закономірностей для значної частини набору даних.

Однак, помітний контраст з'являється при оцінці продуктивності на іншій половині датчиків. На жаль, запропоновані підходи показали неоптимальні результати на цій підмножині, що свідчить про явне обмеження. Ця розбіжність означає, що універсальна архітектура недосяжна для всебічного аналізу і прогнозування всіх датчиків у наборі даних.

Така роздвоєність підкреслює нюанси в характері даних, припускаючи потенційну варіативність або відмінні характеристики між різними типами датчиків. Отже, дослідження вказує на необхідність вивчення адаптованих або

спеціалізованих моделей для підмножин датчиків, визнаючи притаманну їм різноманітність у наборі даних. Отримані результати спонукають до подальшого дослідження основних факторів, що сприяють різним характеристикам, прокладаючи шлях до більш цілеспрямованої розробки та вдосконалення моделей у майбутніх дослідженнях.

3.4. Керівництво користувача

Щоб зручно користуватися розробленою системою для отримання персоналізованих рекомендацій щодо експлуатації техніки, було розроблено клієнт-серверну архітектуру, у вигляді веб-сайту. Вигляд інтерфейсу користувача є мінімалістичний, складається із кнопки вибору категорії, кнопки відправлення даних, та декілька кнопок-картинок для вибору техніки та ступеня складності роботи (див. Рис. 2.3).

При наведенні курсору миші на кнопки вони динамічно змінюють колір, покращуючи користувацький досвід, вказуючи на інтерактивні елементи. Крім того, кнопки на картках змінюються після натискання на них, надаючи користувачам негайний візуальний зворотній зв'язок про зроблений вибір.

Робочий процес системи ретельно продуманий, щоб забезпечити безперебійну роботу користувачів. Для того, щоб продовжити, користувачі повинні вибрати доступну техніку та вказати рівень складності запланованої діяльності. Потім ці вибрані параметри передаються на сервер, використовуючи попередні показники датчиків, що зберігаються в базі даних.

Варто зазначити, що система застосовує логічне обмеження - прогрес зупиняється доти, доки користувачі не обере техніку та складність запланованого завдання. Така навмисна пауза в робочому процесі гарантує, що система отримує необхідні вхідні дані для точного аналізу та генерації рекомендацій.

За допомогою моделей машинного навчання генерують персоналізовані рекомендації на основі вказаних параметрів. Динамічні зміни кольорів в інтерфейсі не тільки сприяють візуальному сприйняттю, але й відіграють функціональну роль, допомагаючи користувачам у процесі вибору. Такий

дизайн, орієнтований на користувача, не тільки покращує естетику, але й сприяє ефективній комунікації та взаємодії з системою, спрощуючи загальний процес прийняття рішень для оптимальної експлуатації важкої техніки.

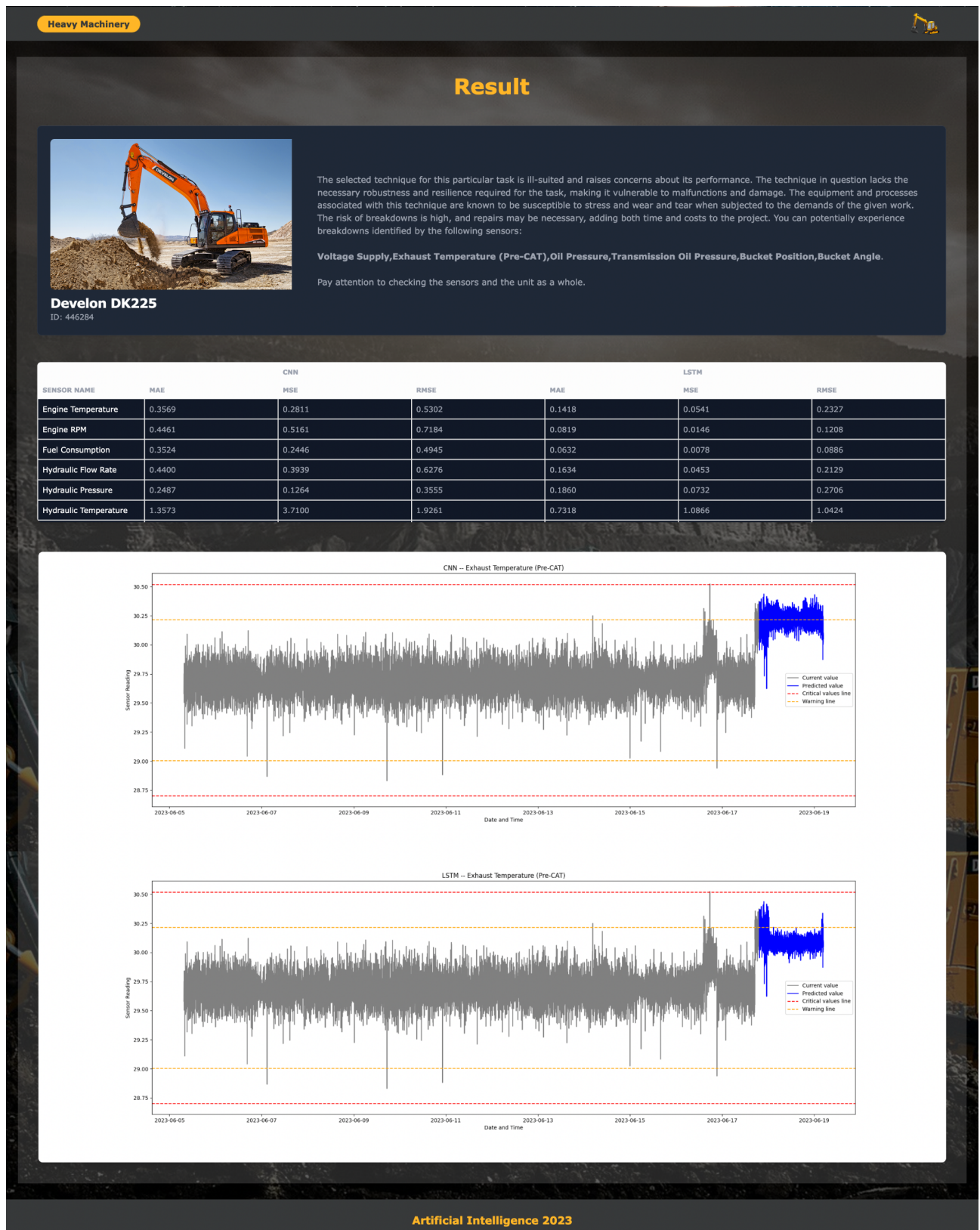


Рис. 3.7. Візуалізації результатів роботи системи

Після передачі даних на сервер виконується прогнозування значень датчиків обладнання. Цей аналітичний процес має на меті спрогнозувати потенційні проблеми в обладнанні, визначаючи датчики, схильні до збоїв або неминучих поломок. Потім сервер складає список виявлених датчиків із відхиленнями, пропонуючи механізм для попередження потенційних несправностей.

Щоб покращити розуміння користувача і полегшити прийняття обґрунтованих рішень, сервер генерує графічні зображення прогнозованих значень датчиків. Ці візуалізації, зображені на рисунку 3.7, слугують потужним інструментом для інтуїтивної інтерпретації даних, надаючи користувачам чіткий і стислий огляд прогнозованих тенденцій.

Візуалізовані результати не тільки представляють прогнозовані значення датчиків, але й містять практичні рекомендації щодо подальшої експлуатації обладнання. Ці рекомендації охоплюють широкий спектр інформації, наприклад, оптимальні режими роботи, потенційний успіх виконання завдань, ймовірність виникнення поломки та пов'язану з нею ймовірність таких подій. Крім того, ймовірність поломки пов'язана з конкретними датчиками, надаючи користувачам цілеспрямовану інформацію про потенційні проблемні зони.

Поєднуючи прогнозну аналітику з інтуїтивно зрозумілою візуалізацією, система надає вичерпний звіт, що дозволяє користувачам приймати обґрунтовані рішення щодо стратегії експлуатації обладнання. Запропонований підхід не лише допомагає у профілактичному обслуговуванні, але й оптимізує продуктивність та довговічність обладнання, проактивно вирішуючи потенційні проблеми до того, як вони загострюються.

3.5. Висновки до розділу

Цей розділ містить детальний опис інструментів розробки, використаних у цьому дослідженні, включаючи обрану мову програмування та пов'язані з нею фреймворки. Обрана мова програмування Python слугує наріжним каменем, а її фреймворки - Flask для бекенду та Angular для користувацького інтерфейсу - сприяють створенню надійної архітектури системи. Крім того, для покращення

аналітичних можливостей системи ретельно використовуються основні дослідницькі бібліотеки, такі як Pandas, NumPy, SciPy, Scikit-Learn, TensorFlow та Keras. В даному розділі також описані тонкощі проектування системи прогнозування для сенсорних даних на основі часових рядів. Систематичне вивчення цих вимог забезпечує всебічне розуміння можливостей системи, гарантуючи, що вона не лише задовольняє потреби користувачів, але й відповідає стандартам продуктивності та експлуатації.

В розділу також описано поетапний аналіз та обробка дослідницьких даних. Виконується ретельний порівняльний аналіз навчених моделей машинного навчання, що супроводжується комплексною оцінкою якості моделей. Ця аналітична база створює надійну основу для оцінки ефективності та надійності прогнозних моделей.

Насамкінець представлено зразковий сценарій роботи системи. Інтерфейс користувача представляє рекомендації щодо правильної експлуатації обладнання на основі запропонованих параметрів. Ця практична ілюстрація підкреслює реальну застосовність і орієнтованість на користувача розробленої системи, демонструючи її потенціал для надання дієвих рекомендацій щодо оптимізації роботи обладнання на основі предиктивної аналітики та машинного навчання.

ВИСНОВКИ

В ході виконання даної магістерської кваліфікаційної роботи по застосуванню машинного навчання щодо оцінки технічних характеристик важкої техніки для формування персоналізованих рекомендацій було зібрано великий набір даних. Був розроблений алгоритм структурування та обробки даних, проаналізовано та натреновано декілька моделей машинного навчання, а саме CNN та LSTM. Для обробки часових рядів було використано метод вікон, для структурування даних перед передаванням їх на вхід моделей.

У ході виконання роботи була реалізована зручна система для користувача, у вигляді веб-сайту. За допомогою неї, користувач може отримати прогнозовані величини по кожному із датчиків та рекомендації щодо використання конкретної техніки на роботах різної складності. Таким чином, мету дослідження було досягнуто.

При виконанні практичної та теоретичної частини роботи були виконані наступні поставлені задачі:

1) Проаналізовано ринок важкої техніки та виявлено низку важливих датчиків для коректного функціонування техніки. Відповідно до аналізу, можна стверджувати, що кожна категорія техніки має індивідуальні важливі параметри, тому було прийнято рішення досліджувати ту частину датчиків, яка притаманна для всіх моделей..

2) Проведено декілька етапів обробки даних для їхньої очистки та кращих результатів прогнозування методів машинного навчання. Серед основних етапів: знаходження та видалення дублікатів даних; видалення аномальних даних; виявлення та обробка відсутніх значень; масштабування даних.

3) Проаналізовано моделі глибинного навчання, створено архітектуру моделі відповідно до типу датчика системи.

4) Розроблено систему для отримання персоналізованих рекомендацій щодо експлуатування важкої техніки. Створено зручний користувацький інтерфейс для вибору важкої техніки та візуалізації результатів.

Підводячи підсумки роботи, слід зауважити, що прогнозування значень датчиків важкої техніки за допомогою аналізу часових рядів має значну актуальність і багатообіцяючі перспективи в різних галузях промисловості. Ось кілька ключових причин, чому це дослідження є актуальним, і потенційні перспективи, які воно відкриває:

- Профілактичне обслуговування та скорочення простоїв: прогностичне моделювання даних датчиків важкої техніки дає змогу передбачити збої та несправності обладнання. Визначаючи закономірності, що вказують на потенційні проблеми, можна планувати технічне обслуговування заздалегідь, мінімізуючи час простою і запобігаючи дорогим поломкам. Такий підхід може суттєво підвищити операційну ефективність та зменшити незаплановані витрати на технічне обслуговування.
- Економія коштів та оптимізація ресурсів: точні прогнози значень датчиків дозволяють організаціям оптимізувати розподіл ресурсів. Розуміючи, коли конкретні компоненти обладнання можуть потребувати уваги або заміни, компанії можуть ефективніше розподіляти ресурси, уникаючи непотрібних замін і мінімізуючи загальні витрати на технічне обслуговування.
- Продовження терміну служби обладнання: своєчасне втручання на основі прогнозних даних може продовжити термін служби важкого обладнання. Вирішуючи потенційні проблеми до того, як вони загострюються, організації можуть зменшити знос, що призведе до продовження терміну служби обладнання та скорочення капітальних витрат на часті заміни.
- Операційна безпека та зменшення ризиків: моніторинг і прогнозування значень датчиків сприяють підвищенню експлуатаційної безпеки. Виявлення аномальних моделей у даних датчиків може запобігти потенційним загрозам безпеці, забезпечуючи безпечне робоче середовище для персоналу. Такий проактивний підхід має вирішальне значення, особливо в галузях, де важке

машинобудування створює невід'ємні ризики.

- Прийняття рішень на основі даних: прогностична аналітика, заснована на даних часових рядів, надає особам, які приймають рішення, цінну інформацію. Організації можуть приймати обґрунтовані рішення щодо розгортання обладнання, стратегій технічного обслуговування та загального операційного планування. Цей процес прийняття рішень на основі даних покращує загальні бізнес-стратегії та конкурентоспроможність.
- Інтеграція з тенденціями Індустрії 4.0 та Інтернету речей: сучасні тенденції Індустрії 4.0 та Інтернету речей (IoT) підкреслюють взаємозв'язок машин і систем. Прогнозна аналітика на основі даних датчиків важкого машинобудування легко узгоджується з цими тенденціями, сприяючи створенню розумних і взаємопов'язаних промислових екосистем. Ця інтеграція може призвести до більш ефективних і автономних промислових процесів.
- Постійне вдосконалення завдяки машинному навчанню: з розвитком машинного навчання прогностичні можливості моделей, що аналізують дані часових рядів, покращуватимуться. Безперервні дослідження в цій галузі відкривають можливості для більш точних і складних алгоритмів, ще більше вдосконалюючи здатність прогнозувати і запобігати проблемам з важкою технікою.

На закінчення, дослідження, спрямовані на прогнозування значень датчиків важкої техніки за допомогою аналізу часових рядів, є актуальними через їх безпосереднє практичне застосування в профілактичному обслуговуванні, економії витрат, підвищенні безпеки та оптимізації ресурсів. Перспективи цього дослідження є багатообіцяючими, оскільки постійний технологічний прогрес сприяє створенню більш досконалих і потужних моделей прогнозування, які можуть докорінно змінити те, як промисловість управляє своїми активами важкого машинобудування.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

- [1] N. Boyko i O. Lukash, «Methodology for Estimating the Cost of Construction Equipment Based on the Analysis of Important Characteristics Using Machine Learning Methods», *J. Eng.*, вип. 2023, с. e8833753, Вер 2023, doi: 10.1155/2023/8833753.
- [2] J. A. Castro-Puma *et al.*, «Automatic learning algorithm for troubleshooting in hydraulic machinery», *Indones. J. Electr. Eng. Comput. Sci.*, вип. 29, вип. 1, с. 535–544, 2023, doi: 10.11591/ijeecs.v29.i1.pp535-544.
- [3] K. Bhushan *et al.*, «Analyzing Reliability and Maintainability of Crawler Dozer BD155 Transmission Failure Using Markov Method and Total Productive Maintenance: A Novel Case Study for Improvement Productivity», *Sustain. Switz.*, вип. 14, вип. 21, 2022, doi: 10.3390/su142114534.
- [4] A. Puška, M. Nedeljković, Ž. Šarkoćević, Z. Golubović, V. Ristić, i I. Stojanović, «Evaluation of Agricultural Machinery Using Multi-Criteria Analysis Methods», *Sustain. Switz.*, вип. 14, вип. 14, 2022, doi: 10.3390/su14148675.
- [5] M. Al-Ghobari, A. Muneer, i S. M. Fati, «Location-aware personalized traveler recommender system (lapta) using collaborative filtering knn», *Comput. Mater. Contin.*, вип. 69, вип. 2, с. 1553–1570, 2021, doi: 10.32604/cmc.2021.016348.
- [6] K. Sarita, R. Devarapalli, S. Kumar, H. Malik, F. P. Garcia Marquez, i P. Rai, «Principal component analysis technique for early fault detection», *J. Intell. Fuzzy Syst.*, вип. 42, вип. 2, с. 861–872, 2022, doi: 10.3233/JIFS-189755.
- [7] P. Aqueveque, L. Radrigan, F. Pastene, A. S. Morales, i E. Guerra, «Data-Driven Condition Monitoring of Mining Mobile Machinery in Non-Stationary Operations Using Wireless Accelerometer Sensor Modules», *IEEE Access*, вип. 9, с. 17365–17381, 2021, doi: 10.1109/ACCESS.2021.3051583.
- [8] S. Kaparathi i D. Bumblauskas, «Designing predictive maintenance systems using decision tree-based machine learning techniques», *Int. J. Qual. Reliab. Manag.*, вип. 37, вип. 4, с. 659–686, Лют 2020, doi: 10.1108/IJQRM-04-2019-0131.
- [9] Y. Bouabdallaoui, Z. Lafhaj, P. Yim, L. Ducoulombier, i B. Bennadji, «Predictive Maintenance in Building Facilities: A Machine Learning-Based Approach», *Sensors*, вип. 21, вип. 4, Art. вип. 4, Січ 2021, doi: 10.3390/s21041044.
- [10] C. Chen, Y. Liu, X. Sun, C. D. Cairano-Gilfedder, i S. Titmus, «An integrated deep learning-based approach for automobile maintenance prediction with GIS data», *Reliab. Eng. Syst. Saf.*, вип. 216, 2021, doi: 10.1016/j.ress.2021.107919.
- [11] S.-L. Lin, «Application of Machine Learning to a Medium Gaussian Support Vector Machine in the Diagnosis of Motor Bearing Faults», *Electronics*, вип. 10, вип. 18, Art. вип. 18, Січ 2021, doi: 10.3390/electronics10182266.
- [12] S. Ayvaz i K. Alpay, «Predictive maintenance system for production lines in

- manufacturing: A machine learning approach using IoT data in real-time», *Expert Syst. Appl.*, вип. 173, 2021, doi: 10.1016/j.eswa.2021.114598.
- [13] K. Velmurugan, P. Venkumar, i P. R. Sudhakara, «SME 4.0: Machine learning framework for real-time machine health monitoring system», *J. Phys. Conf. Ser.*, вип. 1911, вип. 1, с. 012026, Трав 2021, doi: 10.1088/1742-6596/1911/1/012026.
- [14] D. She i M. Jia, «A BiGRU method for remaining useful life prediction of machinery», *Meas. J. Int. Meas. Confed.*, вип. 167, 2021, doi: 10.1016/j.measurement.2020.108277.
- [15] A. Shehadeh, O. Alshboul, R. E. Al Mamlook, i O. Hamedat, «Machine learning models for predicting the residual value of heavy construction equipment: An evaluation of modified decision tree, LightGBM, and XGBoost regression», *Autom. Constr.*, вип. 129, 2021, doi: 10.1016/j.autcon.2021.103827.
- [16] S. Milano, M. Taddeo, i L. Floridi, «Recommender systems and their ethical challenges», *AI Soc.*, вип. 35, вип. 4, с. 957–967, Груд 2020, doi: 10.1007/s00146-020-00950-y.
- [17] M. Sohaib, S. Mushtaq, i J. Uddin, «Deep Learning for Data-Driven Predictive Maintenance», в *Vision, Sensing and Analytics: Integrative Approaches*, вип. 207, М. А. R. Ahad i A. Inoue, Ред., в Intelligent Systems Reference Library, vol. 207. , Cham: Springer International Publishing, 2021, с. 71–95. doi: 10.1007/978-3-030-75490-7_3.
- [18] J. F. Torres, D. Hadjout, A. Sebaa, F. Martínez-Álvarez, i A. Troncoso, «Deep Learning for Time Series Forecasting: A Survey», *Big Data*, вип. 9, вип. 1, с. 3–21, Лют 2021, doi: 10.1089/big.2020.0159.
- [19] H. Widiputra, A. Mailangkay, i E. Gautama, «Multivariate CNN-LSTM Model for Multiple Parallel Financial Time-Series Prediction», *Complexity*, вип. 2021, с. e9903518, Жов 2021, doi: 10.1155/2021/9903518.

Додаток А.

Таблиця А.1. Значення похибок моделей CNN та LSTM

SENSORS	MAE		MSE		RMSE	
	CNN	LSTM	CNN	LSTM	CNN	LSTM
Engine Temperature	0,1900	0,1819	0,1077	0,0630	0,3281	0,2510
Engine RPM	0,4632	0,2438	0,4745	0,1660	0,6888	0,4074
Fuel Consumption	0,4059	0,1747	0,3042	0,0593	0,5515	0,2436
Hydraulic Flow Rate	0,5147	0,2492	0,4997	0,1396	0,7069	0,3737
Hydraulic Pressure	0,2513	0,2365	0,1033	0,0914	0,3214	0,3022
Hydraulic Temperature	1,7879	1,5659	5,0747	3,8935	2,2527	1,9732
Voltage Supply	0,3778	0,3057	0,2725	0,1800	0,5185	0,4242
Current Draw	0,2661	0,1807	0,1544	0,0834	0,3930	0,2887
Battery Voltage	0,2449	0,2317	0,1368	0,1065	0,3699	0,3263
Brake Pressure	0,1699	0,1296	0,0643	0,0399	0,2535	0,1998
Transmission Temperature	1,7343	1,6803	4,8558	4,6231	2,2036	2,1501
Brake Temperature	1,6123	1,4781	4,3578	3,6630	2,0875	1,9139
Coolant Flow Rate	2,0299	1,7748	8,1565	6,1418	2,8560	2,4783
Emission Levels (e.g., CO, NOx, SOx)	0,1012	0,1260	0,0582	0,0595	0,2413	0,2440
Fuel Pressure	0,1498	0,2091	0,0715	0,1026	0,2673	0,3203
Lubricant Pressure	0,3658	0,3283	0,2584	0,2220	0,5083	0,4711
Lubricant Temperature	0,7301	0,7213	1,0308	0,9625	1,0153	0,9810
Radiator Temperature	0,0883	0,0803	0,0136	0,0119	0,1167	0,1092
Exhaust Temperature (Post-CAT)	0,1097	0,0957	0,1197	0,0787	0,3460	0,2806
Exhaust Temperature (Pre-CAT)	0,1182	0,2047	0,0227	0,0549	0,1507	0,2343
Exhaust Gas Flow Rate	0,1161	0,1476	0,0316	0,0417	0,1777	0,2041
Air Intake Temperature	0,1034	0,0927	0,0642	0,0592	0,2534	0,2432
Air Flow Rate	0,2198	0,3150	0,2255	0,2335	0,4749	0,4832
Steering Pressure	0,1526	0,1463	0,0822	0,0751	0,2867	0,2740
Oil Temperature	0,3672	0,2356	0,5676	0,2615	0,7534	0,5114
Oil Pressure	3,1885	2,6682	26,4364	19,7890	5,1416	4,4485
Air Pressure	1,4108	1,0421	11,3669	5,9593	3,3715	2,4412
Torque (Engine)	0,6305	0,6308	0,7607	0,7272	0,8722	0,8528
Torque (Drive Train)	0,7953	0,8672	1,2102	1,2677	1,1001	1,1259
Fluid hydraulic Rate	1,1852	1,1017	2,4711	2,1304	1,5720	1,4596
Transmission Oil Pressure	0,7332	0,5802	2,1574	1,1054	1,4688	1,0514
Transmission Oil Temperature	0,7581	0,7095	1,1792	0,9775	1,0859	0,9887

Exhaust Gas Temperature	1,3391	1,2220	4,5202	3,3242	2,1261	1,8232
Pump Pressure	1,4967	1,3925	3,5430	3,1527	1,8823	1,7756
Oxygen Content	4,2532	3,8505	35,6195	31,0941	5,9682	5,5762
Fuel Flow Rate (for accurate consumption monitoring)	5,5785	5,0802	48,5563	39,8330	6,9682	6,3113
Cabin Temperature	1,2550	0,9652	2,7513	1,6544	1,6587	1,2863
Vibration Level	2,9115	2,5939	14,0630	11,2215	3,7501	3,3499
Load Weight	4,1156	3,7786	31,3512	25,0559	5,5992	5,0056
Noise Frequency Spectrum	4,2534	4,1279	33,2183	29,9020	5,7635	5,3761
Particle Count (e.g., for air quality or filtration monitoring)	4,6517	4,3296	37,3992	33,2349	6,1155	5,7650
Carbon Dioxide (CO2) Content	3,4398	3,0755	22,3162	19,0302	4,7240	4,3624
Bucket Position	6,4721	5,8631	71,8733	60,1740	8,4778	7,7572
Bucket Angle	5,2877	4,8992	48,3343	42,5971	6,9523	6,5266
Noise Level	2,8057	2,2057	14,1577	11,1577	3,7627	2,7627

Додаток Б.

Реалізована система: <https://github.com/LukashOleksii/ng-ml-timeseries-prediction>.