

**ДЕРЖАВНИЙ УНІВЕРСИТЕТ
ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ
НАВЧАЛЬНО-НАУКОВИЙ ІНСТИТУТ
ІНФОРМАЦІЙНИХ ТЕХНОЛОГІЙ
КАФЕДРА ІНЖЕНЕРІЇ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ**

КВАЛІФІКАЦІЙНА РОБОТА

на тему: «Методика аналізу та прогнозування ризиків у
стартап-проєктах на основі технології машинного навчання»

на здобуття освітнього ступеня магістра
зі спеціальності 121 Інженерія програмного забезпечення
освітньо-професійної програми «Інженерія програмного забезпечення»

*Кваліфікаційна робота містить результати власних досліджень. Використання
ідей, результатів і текстів інших авторів мають посилання
на відповідне джерело*

_____ Дмитро ДЗЯДЕВИЧ
(підпис)

Виконав: здобувач вищої освіти групи ПДМ-62
_____ Дмитро ДЗЯДЕВИЧ

Керівник: _____ Володимир САДОВЕНКО
к.ф.-м.н., доц.

Рецензент: _____

**ДЕРЖАВНИЙ УНІВЕРСИТЕТ
ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ**

Навчально-науковий інститут інформаційних технологій

Кафедра Інженерії програмного забезпечення

Ступінь вищої освіти Магістр

Спеціальність 121 Інженерія програмного забезпечення

Освітньо-професійна програма «Інженерія програмного забезпечення»

ЗАТВЕРДЖУЮ

Завідувач кафедри

Інженерії програмного забезпечення

_____ Ірина ЗАМРІЙ

«_____» _____ 2024 р.

**ЗАВДАННЯ
НА КВАЛІФІКАЦІЙНУ РОБОТУ**

_____ Дзядевич Дмитро Дмитрович

1. Тема кваліфікаційної роботи: «Методика аналізу та прогнозування ризиків у стартап-проектах на основі технології машинного навчання»

керівник кваліфікаційної роботи Володимир САДОВЕНКО, к.ф.-м.н., доцент, професор кафедри ІІЗ,

затверджені наказом Державного університету інформаційно-комунікаційних технологій від «15» жовтня 2024 р. № 320.

2. Строк подання кваліфікаційної роботи «17» грудня 2024 р.

3. Вихідні дані до кваліфікаційної роботи: науково-технічна література, параметри, метод прогнозування ризиків, вимоги до точност.

4. Зміст розрахунково-пояснювальної записки (перелік питань, які потрібно розробити)

1. Дослідження методів прогнозування ризиків у стартапах..

2. Аналіз технологій прогнозування ризиків.

3. Розробка та валідація методу прогнозування ризиків у стартап проектах

5. Перелік ілюстративного матеріалу: *презентація*

1. Мета, об'єкта та предмет дослідження

2. Актуальність роботи

3. Постановка задачі

4. Об'єкт кластеризації

5. Алгоритм прогнозування ризиків

6. Дата видачі завдання «16» жовтня 2024 р.

КАЛЕНДАРНИЙ ПЛАН

№ з/п	Назва етапів кваліфікаційної роботи	Строк виконання етапів роботи	Примітка
1	Аналіз наявної науково-технічної літератури	18.10-23.10.2024	
2	Вивчення матеріалів для аналізу розвитку у стартап проєктах	12.11-14.11.2024	
3	Дослідження методів прогнозування ризиків	18.11-20.11.2024	
4	Аналіз особливостей впливу методів прогнозування ризиків	22.11-25.11.2024	
5	Дослідження технологій прогнозування ризиків	27.11-03.12.2024	
6	Застосування прогнозування ризиків у стартап-проєктах на основі технології машинного навчання	07.12-10.12.2024	
7	Оформлення роботи: вступ, висновки, реферат	12.12-20.12.2024	
8	Розробка демонстраційних матеріалів	20.12-26.12.2024	
9	Попередній захист роботи	27.11-17.12.2024	

Здобувач вищої освіти

(підпис)

Дмитро ДЗЯДЕВИЧ

Керівник

кваліфікаційної роботи

(підпис)

Володимир СОДОВЕНКО

РЕФЕРАТ

Текстова частина кваліфікаційної роботи на здобуття освітнього ступеня магістра: 83 стор., 18 табл., 1 рис., 15 фор., 31 джерел.

Мета роботи – оптимізація процесу аналізу і прогнозування ризиків у стартап-проектах за рахунок використання методів машинного навчання, що забезпечують ефективну оцінку рівня ризику та надання рекомендацій для їх мінімізації.

Об'єкт дослідження – процес аналізу і прогнозування ризиків у стартап-проектах.

Предмет дослідження – методи аналізу і прогнозування ризиків у стартап-проектах.

У роботі використано сучасні методи машинного навчання, технології об'єктно-орієнтованого програмування, а також апарат математичної статистики. Основний акцент зроблено на використанні методів глибокого навчання для аналізу та прогнозування ризиків у стартап-проектах.

Проведено огляд та аналіз сучасних підходів до аналізу ризиків, включаючи регресійні моделі, алгоритми підтримки векторів (SVM), ансамблеві методи (Random Forest, Gradient Boosting) та архітектури нейронних мереж (CNN, RNN, трансформери). Розглянуто перспективи інтеграції байєсових методів і Explainable AI (LIME, SHAP) для підвищення прозорості та інтерпретації результатів.

Розроблено багаторівневий метод прогнозування ризиків у стартап-проектах, що враховує динамічні фактори середовища, такі як ринкові умови, структура команди, фінансові показники та інноваційність. Застосовано методи попередньої обробки даних, нормалізації, зменшення розмірності (PCA), а також побудови прогнозних моделей з використанням Random Forest та нейронних мереж.

Реалізовано програмну систему для прогнозування ризиків, розроблену на основі технологій Python та бібліотек NumPy, Pandas, scikit-learn. Програмна система забезпечує інтеграцію різних методів прогнозування, адаптивність до змін середовища та високу точність класифікації стартапів за рівнем ризику.

Отримані результати демонструють перспективність розробленого методу для застосування у сфері інвестування, управління ризиками та стратегічного планування. Метод дозволяє не лише класифікувати стартапи за рівнем ризику, але й надавати рекомендації щодо зменшення ризиків, що робить його ефективним інструментом підтримки прийняття рішень

КЛЮЧОВІ СЛОВА: ПРОГНОЗУВАННЯ РИЗИКІВ, СТАРТАП-ПРОЄКТИ, МАШИННЕ НАВЧАННЯ, ГЛИБОКЕ НАВЧАННЯ, RANDOM FOREST, SVM, НЕЙРОННІ МЕРЕЖІ, АНАЛІЗ ДАНИХ.

ABSTRACT

The text part of the qualifying work for obtaining a master's degree: 83 pages, 18 tables, 1 figure, 15 figures, 31 sources.

The purpose of the work is to optimize the process of risk analysis and forecasting in startup projects through the use of machine learning methods that provide effective assessment of the level of risk and the provision of recommendations for their minimization.

The object of the study is the process of risk analysis and forecasting in startup projects.

The subject of the study is methods of risk analysis and forecasting in startup projects.

The work uses modern methods of machine learning, object-oriented programming technologies, as well as the apparatus of mathematical statistics. The main emphasis is on the use of deep learning methods for risk analysis and forecasting in startup projects.

A review and analysis of modern approaches to risk analysis, including regression models, support vector algorithms (SVM), ensemble methods (Random Forest, Gradient Boosting) and neural network architectures (CNN, RNN, transformers) were conducted. Prospects for integration of Bayesian methods and Explainable AI (LIME, SHAP) to increase transparency and interpretation of results are considered.

A multi-level method of risk forecasting in startup projects has been developed, which takes into account dynamic environmental factors such as market conditions, team structure, financial indicators and innovativeness. The methods of data preprocessing, normalization, dimensionality reduction (PCA), as well as construction of predictive models using Random Forest and neural networks were applied.

A software system for risk forecasting developed on the basis of Python technologies and NumPy, Pandas, scikit-learn libraries has been implemented. The software system provides integration of various forecasting methods, adaptability to environmental changes and high accuracy of classification of startups by risk level.

The obtained results demonstrate the prospects of the developed method for application in the field of investment, risk management and strategic planning. The method allows not only to classify startups by risk level, but also to provide recommendations for reducing risks, which makes it an effective decision support tool

KEY WORDS: RISK PREDICTION, STARTUP PROJECTS, MACHINE LEARNING, DEEP LEARNING, RANDOM FOREST, SVM, NEURAL NETWORKS, DATA ANALYSIS.

ЗМІСТ

ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ	11
ВСТУП.....	12
1 АНАЛІЗ ПРЕДМЕТНОЇ ГАЛУЗІ	14
1.1 Стартап-проекти та їх специфіка.....	14
1.2 Ризики у стартапах.....	16
1.3 Огляд літератури та існуючих підходів до аналізу ризиків у стартапах	19
2 МЕТОДИКА АНАЛІЗУ ТА ПРОГНАЗУВАННЯ РИЗИКІВ НА ОСНОВІ МАШИНОГО НАВЧАННЯ	29
2.1 Аналіз існуючих методик прогнозування ризиків на основі машинного навчання	29
2.1.1 Класичні методи аналізу та прогнозування ризиків	29
2.2 Сучасні методи прогнозування ризиків з використанням машинного навчання	53
3 РОЗРОБКА МЕТОДУ АНАЛІЗУ ТА ПРАГНАЗУВАННЯ РИЗИКІВ У СТАРТАП -ПРОЄКТАХ НА ОСНОВІ ТЕХНОЛОГІЇ МАШИНОГО	60
3.1 Опис розробки методу	60
3.2 Опис роботи алгоритма	63
3.3 Опис використаних програмних засобів.....	67
ВИСНОВКИ.....	72
ПЕРЕЛІК ПОСИЛАНЬ	75
ДОДАТОК А. ДЕМОНСТРАЦІЙНІ МАТЕРІАЛИ	77
ДОДАТОК Б. ЛІСТИНГИ ПРОГРАМНИХ МОДУЛІВ.....	83

ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ

CRM - Управління відносинами з клієнтами (*customer relationship management*).

MVP - Мінімально життєздатний продукт (Minimum viable product).

GDPR – Загальний регламент про захист даних (General Data Protection Regulation).

SXSW - З півдня на південний захід (South by Southwest).

МНК - Метод найменших квадратів.

GLM - Библиотека для OpenGL (OpenGL Mathematics).

SVM - Метод опорних векторів (support vector machine).

RNN - Рекуррентная нейронная сеть (Recurrent neural network)

LSTM - Довга короткострокова пам'ять (Long short-term memory).

PCA - Метод основних компонент

ВСТУП

З розвитком сучасних технологій та зростаючою динамікою ринку постає необхідність у вдосконаленні методів управління ризиками, особливо у стартап-проектах. Стартапи функціонують у середовищі високої невизначеності, де обмежені ресурси, швидкі організаційні зміни та нестабільні ринкові умови створюють численні виклики для їхньої успішної реалізації.

Сучасні підходи до аналізу ризиків у стартапах значною мірою спираються на технології машинного навчання, що дають змогу обробляти великі обсяги даних, виявляти приховані закономірності та формувати прогнози. Використання таких методів, як нейронні мережі, алгоритми підтримки векторів (SVM), ансамблеві моделі (Random Forest, Gradient Boosting) та методи Explainable AI (LIME, SHAP), дозволяє значно підвищити точність і стабільність прогнозів.

Даний диплом присвячений розробці багаторівневого методу прогнозування ризиків у стартап-проектах із використанням сучасних технологій машинного навчання. Проведено аналіз існуючих підходів до управління ризиками, визначено їх переваги та недоліки, а також запропоновано нову методологію, яка інтегрує байєсові принципи, алгоритми машинного навчання та механізми інтерпретації.

Ця робота спрямована на вирішення актуальних проблем управління ризиками, пов'язаних з високою мінливістю ринку, обмеженістю ресурсів та складністю багатofакторних процесів у стартапах. Запропонований підхід має на меті забезпечити надійні прогнози, які допоможуть підприємцям, інвесторам та аналітикам приймати обґрунтовані рішення.

Таким чином, завдання розробки інноваційного методу прогнозування ризиків у стартап-проектах із використанням машинного навчання є сучасним та актуальним.

Мета роботи – оптимізація процесу аналізу і прогнозування ризиків у стартап-проектах за рахунок використання методів машинного навчання, що забезпечують ефективну оцінку рівня ризику та надання рекомендацій для їх мінімізації.

Об'єкт дослідження – процес аналізу і прогнозування ризиків у стартап-проектах.

Предмет дослідження – машинного навчання та алгоритми кластеризації, які використовуються для аналізу і прогнозування ризиків у стартапах.

Для досягнення мети вирішувалися наступні завдання:

1. Літературний огляд. Проведено аналіз сучасних наукових джерел для вивчення існуючих підходів до прогнозування ризиків та технологій машинного навчання.
2. Аналіз даних. Виконано попередню обробку даних, включаючи нормалізацію, зменшення розмірності та виявлення ключових факторів ризику.
3. Моделювання. Розроблено багаторівневу методологію прогнозування ризиків, що включає інтеграцію байєсових принципів, алгоритмів SVM, Random Forest та нейронних мереж.
4. Реалізація. Створено програмну систему для прогнозування ризиків із використанням Python та відповідних бібліотек (NumPy, Pandas, scikit-learn).
5. Валідація. Проведено експерименти для оцінки точності та стабільності розробленої моделі.
6. Оптимізація. Удосконалено метод для підвищення точності, швидкості та адаптивності прогнозування.

Вирішення цих завдань сприяє створенню надійної системи прогнозування ризиків, яка може знайти застосування у сфері інвестування, управління проектами та стратегічного планування.

1 АНАЛІЗ ПРЕДМЕТНОЇ ГАЛУЗІ

1.1 Стартап-проекти та їх специфіка

Стартапи — це новаторські підприємства, які створюють інноваційні продукти, послуги або технології з великим потенціалом для швидкого розвитку. Їхня основна мета полягає не лише у виході на ринок, а й у завоюванні значної його частки, встановленні нових стандартів чи навіть зміні правил гри в галузі. Відмінною особливістю стартапів є їхня орієнтація на інновації, швидке зростання та гнучкість у прийнятті рішень. Однак ці риси супроводжуються високим рівнем невизначеності, складнощами з фінансуванням, обмеженими ресурсами та іншими викликами, які суттєво відрізняють їх від традиційних бізнесів. Саме специфіка стартапів формує необхідність у спеціалізованих підходах до управління, аналізу ринкових умов і мінімізації ризиків.

Одним із найважливіших аспектів специфіки стартапів є їхня інноваційність. Завдяки здатності впроваджувати нові ідеї та технології вони можуть вирішувати існуючі проблеми ефективніше або відкривати нові ринки. Інноваційність робить такі проекти привабливими для інвесторів і споживачів, але водночас збільшує рівень ризиків. Розробка радикально нових продуктів чи послуг, як-от штучний інтелект для автоматизації процесів або мобільні платформи для спільного використання ресурсів, вимагає значних ресурсів та глибокого розуміння ринкових умов.

Продуктові інновації, такі як рішення на основі штучного інтелекту від OpenAI або моделі спільного використання ресурсів, створені Uber і Airbnb, наочно демонструють, як інноваційні підходи можуть трансформувати цілі галузі. Водночас такі рішення часто стикаються з викликами: недостатньою готовністю ринку, технічними обмеженнями та високими витратами на реалізацію.

Ще одним важливим аспектом специфіки стартапів є їхня здатність адаптувати традиційні бізнес-моделі або створювати нові. Заміна разового

продажу підписною моделлю, як це зроблено в Netflix, або freemium-підходи, популяризовані Spotify, дають змогу залучати широку аудиторію, одночасно забезпечуючи стабільний потік доходів. Але такі бізнес-моделі потребують точного планування й розрахунків, адже без чіткого фінансового контролю масштабування може призвести до збитків. Процесні інновації також відіграють важливу роль.

Наприклад, впровадження хмарних технологій або CRM-систем допомагає оптимізувати операційні витрати, проте такі рішення можуть бути недоступними для невеликих стартапів через обмежені фінансові ресурси.

Середовище, у якому працюють стартапи, вирізняється високим рівнем невизначеності. Вони часто опиняються на ринках, які ще не сформовані, і працюють у сферах, де поведінка споживачів може бути непередбачуваною. Це створює значні ризики, але водночас стимулює розвиток адаптивних підходів до планування та аналізу. Наприклад, стартапи, що впроваджують екологічні рішення, стикаються з проблемою високих витрат і недостатньою готовністю клієнтів платити за інноваційні продукти.

Навіть успішні проекти, такі як iPhone, на початкових етапах могли зіштовхнутися з нерозумінням користувачів. Додатковий тиск створює конкуренція, адже великі корпорації з більшими ресурсами здатні швидко адаптувати ідеї та вийти на ринок з аналогічними продуктами.

Дефіцит ресурсів є ще однією характерною рисою стартапів. Більшість із них залежить від зовнішніх джерел фінансування, таких як венчурні інвестиції або гранти. Проте високий рівень конкуренції за інвестиції робить цей процес складним і часом непрогнозованим.

Невеликі команди часто виконують кілька функцій одночасно, що підвищує ризик вигорання та знижує ефективність. Крім того, брак доступу до сучасних технологій може уповільнити розробку продуктів або навіть зірвати плани масштабування.

Динамічне середовище, у якому працюють стартапи, вимагає постійної готовності до змін. Технологічний прогрес змушує їх постійно вдосконалювати

свої продукти та сервіси, щоб залишатися конкурентоспроможними. Водночас зміна споживчих звичок і вплив соціальних трендів, таких як екологічність чи інклюзивність, можуть суттєво змінити попит на їхні рішення. Здатність адаптуватися та впроваджувати зміни в таких умовах є критичною для виживання та розвитку.

Командна робота також є одним із ключових факторів успіху стартапів. Компетентність, згуртованість і здатність співробітників швидко приймати рішення визначають, наскільки ефективно проєкт зможе подолати виклики. У разі втрати ключових працівників чи виникнення внутрішніх конфліктів робота може бути паралізованою, а плани — зірваними. Для уникнення цих ризиків важливо створювати прозорі процеси управління, ефективну комунікацію та систему мотивації співробітників.

1.2 Ризики у стартапах

Ризики у стартапах є важливою і невід’ємною частиною їхнього існування, оскільки саме вони формують виклики, що супроводжують розвиток інноваційних проєктів. Для нових компаній, які працюють у нестабільному середовищі з обмеженими ресурсами та високим рівнем конкуренції, управління ризиками стає критично важливим інструментом для досягнення успіху.

Високий рівень невизначеності, властивий новаторським проєктам, обумовлює потребу в систематичному аналізі та ефективному підході до управління ризиками. Відсутність таких механізмів може стати причиною краху навіть найбільш перспективних ідей.

Фінансові ризики є найпоширенішими і водночас найбільш руйнівними для стартапів. Через те, що більшість нових компаній залежать від зовнішніх джерел фінансування, таких як венчурний капітал, бізнес-ангели або гранти, будь-які невдачі у фінансовому плануванні можуть призвести до катастрофічних наслідків. Наприклад, залучення фінансування є складним процесом, який потребує чіткої презентації проєкту, доказів його перспективності та довіри інвесторів.

Навіть за умови успішного старту проєкту неправильне управління бюджетом може призвести до виснаження коштів ще до виходу продукту на ринок. Такі фактори, як затримки у розробці, зростання витрат на технології або неочікувані витрати на маркетинг, лише підсилюють цей ризик.

Для зменшення фінансових ризиків важливо розробити чіткий фінансовий план, який враховує всі можливі сценарії, і забезпечити регулярний контроль за витратами.

Ринкові ризики є ще однією значною перешкодою на шляху до успіху стартапу. Невизначеність попиту, непередбачувана поведінка споживачів та високий рівень конкуренції створюють серйозні виклики. Багато інноваційних ідей, які здавалися перспективними на етапі розробки, виявилися незатребуваними через те, що ринок не був готовий до їхнього впровадження. Наприклад, продукти, які змінюють звички користувачів, потребують значних інвестицій у маркетинг для підвищення обізнаності.

Водночас великі корпорації з більшими ресурсами можуть швидко адаптувати новаторські ідеї і витіснити стартапи з ринку. Для мінімізації цих ризиків стартапам необхідно проводити ретельний аналіз ринку, тестувати мінімально життєздатний продукт (MVP) і залучати зворотний зв'язок від потенційних клієнтів.

Операційні ризики є невід'ємною складовою стартапів через їхню обмежену організаційну структуру. Проблеми можуть виникати через недостатню координацію роботи команди, конфлікти між засновниками, втрату ключових співробітників або неправильне управління проєктами.

У невеликих командах, які часто характерні для стартапів, кожен член виконує декілька функцій, що збільшує ризик вигорання та помилок через перевантаження. Крім того, відсутність чіткої стратегії або неадекватно поставлені цілі можуть уповільнити розвиток проєкту або призвести до марнотратства ресурсів. Для зменшення операційних ризиків стартапи мають забезпечити прозорість процесів управління, налагодження комунікації між співробітниками та навчання членів команди.

Технологічні ризики особливо актуальні для стартапів, що орієнтуються на інноваційні рішення. Висока швидкість технологічного прогресу створює як можливості, так і виклики. Наприклад, розробка складних технологій, таких як штучний інтелект чи блокчейн, потребує значних інвестицій у дослідження, що може стати непосильним для стартапу.

Навіть успішний продукт може швидко застаріти через появу нових технологій або конкурентів із кращими рішеннями. Для зменшення цих ризиків стартапи мають інвестувати у постійне вдосконалення своїх продуктів, захищати свої розробки патентами і впроваджувати адаптивні технологічні стратегії.

Регуляторні ризики часто недооцінюються, але вони можуть мати значний вплив на діяльність стартапу. Зміни в законодавстві, особливо у сферах, що стосуються захисту даних, фінансів або охорони здоров'я, можуть створити додаткові бар'єри для розвитку. Наприклад, впровадження Загального регламенту захисту даних (GDPR) у Європейському Союзі вимагало від багатьох стартапів перегляду їхніх бізнес-моделей, що потребувало додаткових витрат. Постійний моніторинг змін у законодавстві та консультації з юридичними експертами є обов'язковими для мінімізації регуляторних ризиків.

Систематичний підхід до управління ризиками є основою для забезпечення стійкого розвитку стартапу. Включення аналізу ризиків у бізнес-планування, створення резервного фонду та регулярний моніторинг ключових показників діяльності дозволяють швидко реагувати на виклики.

Наприклад, фінансові ризики можна мінімізувати через диверсифікацію джерел фінансування, а ринкові — завдяки постійному аналізу потреб клієнтів. Успішне управління операційними ризиками вимагає розробки ефективних внутрішніх процесів, а технологічні ризики зменшуються шляхом інвестицій у дослідження та захисту інтелектуальної власності.

Розуміння ризиків у стартапах є важливою умовою для їхнього успішного розвитку. Хоча ризики є невід'ємною частиною роботи з інноваціями, вони також стимулюють пошук нестандартних рішень і підвищення ефективності управління. Ефективний ризик-менеджмент дозволяє стартапам зберігати

конкурентоспроможність і досягати довгострокового успіху навіть у найскладніших умовах.

1.3 Огляд літератури та існуючих підходів до аналізу ризиків у стартапах

У сучасному світі стартапи стикаються з численними викликами, які виникають через невизначеність ринку, швидку зміну технологій і високий рівень конкуренції. Для успішного розвитку в цих умовах необхідно розуміти, як аналізувати ризики, розробляти стратегії та адаптуватися до змін. Дослідження провідних авторів, таких як Brad Feld, Eric Rees, Daniel Kahneman, Clayton Christensen та Adam Grant, пропонують цінні інструменти й підходи, які допомагають підприємцям вирішувати ці завдання. Їхні праці охоплюють фінансові, когнітивні, інноваційні та організаційні аспекти управління стартапами, що робить їх незамінними для створення стійких бізнес-моделей.

Eric Ries — автор книги «The Lean Startup: How Today's Entrepreneurs Use Continuous Innovation to Create Radically Successful Businesses»[1], яка стала однією з найвпливовіших робіт у сфері підприємництва. Визнаний експерт у галузі стартап-менеджменту, він розробив концепцію «бережливого стартапу», яка акцентує увагу на швидкості, гнучкості та ефективному використанні ресурсів.

Ries — засновник кількох технологічних стартапів і частий спікер на міжнародних конференціях, таких як SXSW, TechCrunch та Lean Startup Week. Його методика стала основою для багатьох успішних стартапів і широко використовується в інноваційних корпораціях..

Основа ідея є розробка мінімально життєздатного продукту (MVP), який дозволяє компанії швидко перевірити свою ідею на ринку з мінімальними витратами.

Ріс пояснює, що MVP дає можливість уникнути великих інвестицій у продукт, який може виявитися незатребуваним.

Центральною ідеєю книги є цикл «Побудуй – Вимірй – Навчися», який дозволяє компаніям швидко тестувати гіпотези, отримувати зворотний зв'язок і вдосконалювати свій продукт.

- Побудуй: Швидке створення MVP.
- Вимірй: Збір даних від користувачів, що дозволяє оцінити, чи відповідає продукт їхнім потребам.
- Навчися: Аналіз даних для прийняття рішень про подальші дії.

Автор акцентує увагу на тому, що стартапи часто працюють із обмеженими ресурсами. Він пропонує:

- Інвестувати лише в ті функції продукту, які мають підтверджену цінність для клієнтів.
- Уникати зайвих витрат на маркетинг до моменту, поки продукт не буде протестований і оптимізований.

Одним із ключових аспектів методології Lean Startup є можливість швидкої адаптації до змін на ринку. Стартапи, які використовують цей підхід, можуть оперативно реагувати на зворотний зв'язок від клієнтів та вносити необхідні корективи в продукт або послугу. Це знижує ризик невдачі, оскільки компанія постійно вчиться та вдосконалюється на основі реальних даних.

Методика автора спрямована на оптимізацію використання ресурсів. Вона дозволяє уникати зайвих витрат, спрямовуючи фінансові та людські ресурси на найбільш критичні елементи продукту. Це особливо важливо для стартапів з обмеженим бюджетом, де ефективне використання коштів може визначити успіх чи невдачу проекту..

Методологія Lean Startup є надзвичайно корисною для аналізу та управління ризиками в стартапах. Вона дозволяє підприємцям:

- Ідентифікувати ринкові ризики ще на етапі розробки продукту: Постійне тестування гіпотез та отримання зворотного зв'язку від потенційних клієнтів допомагає виявити можливі проблеми та ризики на ранніх стадіях.
- Швидко адаптуватися до зворотного зв'язку від клієнтів: Гнучкість

методології дозволяє вносити зміни в продукт або стратегію практично в режимі реального часу, що підвищує шанси на задоволення потреб ринку.

- Ефективно використовувати обмежені ресурси, мінімізуючи фінансові втрати: Зосереджуючись на найважливіших аспектах та уникаючи непотрібних витрат, стартапи можуть досягати своїх цілей з меншими ризиками.

Daniel Kahneman — видатний американсько-ізраїльський психолог і лауреат Нобелівської премії з економіки 2002 року. Його дослідження в галузі поведінкової економіки та психології прийняття рішень радикально змінили наше розуміння того, як люди сприймають ризики та ухвалюють рішення в умовах невизначеності. Його книга «Thinking is fast and slow»[2] стала світовим бестселером і є фундаментальною працею в аналізі когнітивних процесів.

Роботи Kahneman мають величезне значення для стартапів, які часто змушені ухвалювати швидкі рішення в стресових умовах. Його теорії допомагають краще зрозуміти власні когнітивні упередження та оцінити поведінку клієнтів, партнерів і інвесторів.

У своїй книзі автор описує модель, за якою людський мозок працює через дві системи мислення: Система 1 — швидке мислення, яке є автоматичним, інтуїтивним та емоційним, і Система 2 — повільне мислення, яке є свідомим, раціональним та аналітичним. Система 1 дозволяє швидко ухвалювати рішення, але схильна до когнітивних упереджень.

У контексті стартапів засновник може миттєво оцінити ринкову можливість, але ця оцінка може бути суб'єктивною та необґрунтованою. Система 2 використовується для вирішення складних проблем і потребує значного часу та зусиль, застосовуючись для аналізу фінансових моделей або розробки довгострокової стратегії.

Баланс між цими двома системами є ключовим для успіху. Надмірна залежність від швидкого мислення може призвести до критичних помилок, тоді як занадто повільне мислення може уповільнити процес прийняття рішень.

Автор детально досліджує когнітивні упередження — автоматичні помилки

мислення, що впливають на сприйняття ризиків і ухвалення рішень. Найбільш релевантними для стартапів є оптимістичне упередження, коли засновники переоцінюють свої шанси на успіх, ігноруючи потенційні перешкоди; ефект підтвердження, що полягає в схильності шукати інформацію, яка підтверджує вже існуючі переконання, і відкидати протилежні дані; та якірний ефект, де перша отримана інформація впливає на подальші судження — наприклад, початкова оцінка компанії може визначити всі наступні інвестиційні переговори.

Разом з Amos Tversky Kahneman розробив теорію перспектив, яка пояснює поведінку людей при прийнятті рішень в умовах невизначеності. Люди схильні уникати ризику при можливості прибутку, обираючи гарантований менший прибуток замість ризикованого, але більшого, але схильні до ризику при можливості втрат, готові ризикувати більше, щоб уникнути втрат. Розуміння цієї теорії допомагає засновникам краще оцінювати ризики та ухвалювати обґрунтовані рішення. У маркетингових стратегіях реклама, що підкреслює потенційні втрати від невикористання продукту, може бути ефективнішою.

Adam Grant, видатний американський психолог, автор бестселерів і професор Уортонської школи бізнесу Пенсильванського університету, зробив значний внесок у розуміння того, як розвивати творчість і стимулювати інновації в бізнесі. Його наукові дослідження охоплюють теми творчості, мотивації, лідерства та організаційної ефективності. Він відомий своїми книгами «Originals: How Non-Conformists Move the World» та «Give and Take: Why Helping Others Drives Our Success», які стали бестселерами і глибоко вплинули на сучасне бізнес-мислення.

У книзі «Originals»[3] автор аналізує, як люди, що мислять нестандартно, створюють проривні ідеї та змінюють правила гри в різних галузях. Він наголошує, що оригінальні ідеї часто виглядають недосконалими на початкових етапах, і наводить приклади компаній на кшталт Google та Warby Parker, які стали лідерами ринку, незважаючи на початковий скепсис.

Автор пояснює, як розпізнати потенціал ідеї, яка на перший погляд може здаватися незвичайною чи ризикованою, і підкреслює, що творчі люди не обов'язково є ризикованими авантюристами. Навпаки, вони часто ретельно зважують свої рішення і балансують між стабільністю та інноваціями.

Автор пропонує практичні стратегії для стимулювання творчості в командах, такі як заохочення різноманітності ідей, провокація обговорення та впровадження культури експериментів. Він рекомендує створювати середовище, де співробітники можуть пробувати нові підходи без страху зазнати невдачі, і підкреслює важливість конструктивної критики, яка сприяє покращенню ідей та виявленню їхніх слабких сторін на ранніх етапах.

Для розвитку креативного мислення в стартапах автор пропонує методики, що допомагають розвивати креативність. Важливо заохочувати нестандартні ідеї, навіть якщо вони здаються сумнівними на перший погляд, і навчатися на невдачах, які є невід'ємною частиною творчого процесу.

Практичне значення для стартапів полягає у впровадженні інструментів для стимулювання творчості та інновацій, таких як створення культури довіри, заохочення співпраці між різними департаментами та експертами, а також експериментування з новими ідеями, навіть якщо вони здаються ризикованими. Дослідження Adam Granta відіграють ключову роль у розумінні того, як стимулювати творчість у динамічному середовищі стартапів.

Його праці допомагають підприємцям і керівникам створювати інноваційні продукти та послуги, розвивати культуру творчості та впроваджувати нестандартні рішення для досягнення конкурентних переваг.

Clayton Christensen видатний американський професор бізнесу та один із найвпливовіших мислителів у сфері інноваційного менеджменту, зробив значний внесок у розуміння того, як інновації впливають на успіх компаній.

Його книга «The Innovator's Dilemma» стала класикою вивчення інноваційного менеджменту та отримала широке визнання серед керівників

компаній, засновників стартапів та дослідників. У цій праці автор вводить поняття "disruptive innovations" — технологій або рішень, що починаються з низькоконкурентних сегментів ринку, але поступово замінюють існуючі продукти або послуги.

Christensen зазначає, що великі компанії часто не помічають потенціалу таких інновацій, оскільки зосереджуються на покращенні своїх існуючих продуктів для задоволення поточних клієнтів. Це призводить до того, що нові гравці, орієнтуючись на "disruptive innovations", можуть захоплювати ринок, розвиваючи свої продукти і стаючи лідерами.

Розривні інновації характеризуються доступністю, оскільки зазвичай є дешевшими та простішими у використанні, ніж традиційні аналоги. Вони націлені на нові сегменти клієнтів, які раніше не могли дозволити собі існуючі рішення, і завдяки ефекту масштабу стають конкурентоспроможними, витісняючи традиційних гравців.

Головна проблема, яку описує автор, — це «The Innovator's Dilemma»: великі компанії, прагнучи задовольнити своїх поточних клієнтів, інвестують у вдосконалення існуючих продуктів, ігноруючи потенційні ринки, створені розривними інноваціями. Це дозволяє стартапам захоплювати нові сегменти, розвивати свої продукти і ставати лідерами.

Для вирішення цієї дилеми пропонуються практичні рекомендації. Рекомендується створювати окремі команди або підрозділи, що працюють над розривними технологіями, щоб уникнути конфлікту з основним бізнесом. Диверсифікація продуктів та інвестування в різні технології допомагають знизити ризики, пов'язані з конкуренцією. Постійне стеження за ринковими трендами та адаптація до змін дозволяють компаніям вчасно реагувати на нові виклики. Підкреслюється важливість фокусу на незадоволених потребах клієнтів, що може стати джерелом інноваційних ідей.

Окрім «The Innovator's Dilemma[4]», Christensen написав низку інших

впливових праць. У книзі «The Innovator's Solution» пропонуються конкретні стратегії для компаній, які прагнуть створювати розривні інновації. Обговорюються методи визначення незадоволених потреб клієнтів, побудови бізнес-моделей, що сприяють розвитку інновацій, та управління ризиками при запуску нових продуктів. У «Competing Against Luck» вводиться концепція "Jobs to Be Done" — ідея про те, що клієнти "наймають" продукти для виконання конкретних завдань.

Стверджується, що успішні інновації — це ті, які вирішують конкретні проблеми або потреби споживачів, і пропонуються методи глибокого аналізу цих потреб.

У книзі «How Will You Measure Your Life?» автор зосереджується на особистісному розвитку та управлінні кар'єрою, що має важливе значення для підприємців, які прагнуть знайти баланс між бізнесом та особистим життям. Обговорюється, як принципи управління та інновацій можуть застосовуватися до особистого життя, допомагаючи досягати як професійного успіху, так і особистого задоволення.

Внесок цього мислителя у стартап-екосистему є значним. Його теорія розривних інновацій пояснює, як молоді компанії можуть конкурувати з великими корпораціями, використовуючи нові технології та бізнес-моделі. Наголошується на важливості розуміння потреб клієнтів для створення успішних продуктів і допомозі підприємцям у розробці стратегій для виходу на нові ринки. Стартапи можуть орієнтуватися на незаповнені ніші, пропонуючи дешевші або більш зручні рішення, такі як FinTech-проекти, що спрощують доступ до фінансових послуг, або освітні платформи з доступом до знань за нижчою ціною..

Інноваційні бізнес-моделі, описані в його працях, дозволяють знизити ризики, пов'язані з запуском нових продуктів. Стартапи можуть використовувати розривні інновації як інструмент для створення конкурентної переваги навіть у галузях, де домінують великі корпорації. Його ідеї стали джерелом натхнення для

багатьох підприємців, допомагаючи розуміти природу інновацій та будувати стратегії успіху.

Загалом автор залишив величезний інтелектуальний спадок, який формує сучасні уявлення про управління інноваціями. Його дослідження мають практичне значення як для великих компаній, так і для стартапів, що прагнуть адаптуватися до нових умов ринку та досягти успіху в конкурентному середовищі. Для підприємців, менеджерів та дослідників його праці є незамінним ресурсом, що допомагає розуміти складність інноваційного процесу та ефективно керувати ним. Його ідеї продовжують впливати на бізнес-стратегії по всьому світу, сприяючи розвитку інновацій та економічного зростання.

Brad Feld, один із провідних венчурних капіталістів світу, має багаторічний досвід у сфері інвестицій у стартапи та розвитку підприємницької екосистеми. Його книга «Venture Deals: Be Smarter Than Your Lawyer and Venture Capitalist» є одним із найбільш визнаних посібників для підприємців, які прагнуть залучити інвестиції та зрозуміти механізми роботи венчурного капіталу..

У своїй праці автор детально розглядає структуру венчурних угод, зокрема такі аспекти, як конвертовані позики, акції з преференціями та опціони для співробітників. Конвертовані позики описуються як один із найбільш популярних інструментів для початкового фінансування стартапів. Цей метод дозволяє уникнути необхідності оцінювати компанію на ранніх стадіях, коли її вартість важко визначити. Замість цього інвестори отримують право конвертувати свої позики в частку компанії під час наступного раунду фінансування. Правильно укладений договір про конвертовані позики знижує ризик юридичних суперечок між засновниками та інвесторами.

Акції з преференціями є ключовим елементом угод із венчурними капіталістами. Вони надають інвесторам пріоритет у розподілі прибутків і захист у разі ліквідації компанії. Наприклад, інвестори можуть мати право повернути свої кошти перед тим, як засновники отримають свою частку. Також детально

розглядаються опціонні програми, які дозволяють стартапам залучати таланти без значних фінансових витрат. Правильно побудована система опціонів стимулює співробітників працювати на довгострокову перспективу.

Підготовка до переговорів з інвесторами є ще однією важливою темою. Автор радить підприємцям залучати кваліфікованих юристів із досвідом роботи у сфері венчурного капіталу та детально пояснює, як правильно укладати терм-шити (*Term Sheets*) — документи, які визначають основні умови майбутньої угоди. Він попереджає про типові юридичні помилки, які можуть коштувати компанії дорого, наприклад, нерівномірний розподіл прав голосу, що може призвести до втрати засновниками контролю над компанією.

Для успішної аргументації під час переговорів рекомендується використовувати дані про ринок та реалістичні прогнози росту. Підприємці повинні чітко демонструвати знання про свою цільову аудиторію та конкурентів, а також надавати переконливі фінансові прогнози, що підтверджують перспективність бізнесу.

Стосунки з інвесторами є критично важливими для успіху стартапу. Побудова довгострокових відносин, регулярна комунікація та прозорість у фінансових звітах допомагають зміцнити довіру між сторонами. Автор наводить приклади, коли правильне управління конфліктами дозволило зберегти відносини між інвесторами та засновниками, підкреслюючи важливість відкритого діалогу, особливо в ситуаціях фінансових труднощів.

Для підприємців надаються інсайти щодо управління фінансами та збереження контролю над компанією. Ефективне фінансове управління включає створення резервних фондів для покриття непередбачених витрат та оптимізацію витрат на маркетинг і розробку продукту. Щоб уникнути втрати контролю при залученні великих інвестицій, автор рекомендує зберігати контрольний пакет акцій і уникати угод, які передбачають надмірні права для інвесторів.

Окрім "Venture Deals", Brad Feld є автором інших впливових книг, таких як «Startup Communities: Building an Entrepreneurial Ecosystem in Your City» та «The Startup Community Way». У цих працях він досліджує створення стартап-екосистем у містах, пояснюючи, як сприяти співпраці між підприємцями, інвесторами, університетами та владою. Він наводить приклади успішних спільнот, зокрема екосистему в Боулдері, штат Колорадо, та аналізує, як міста можуть стати центрами інновацій та залучення інвестицій.

Brad Feld також активно ділиться своїми думками через популярний блог, де обговорює теми підприємництва, інвестицій та розвитку технологій. Його статті охоплюють широкий спектр питань — від ризик-менеджменту до побудови команд у стартапах.

Праці цього венчурного капіталіста стали ключовими для багатьох підприємців, які прагнуть глибше зрозуміти механізми венчурного капіталу. Його книги допомогли численним стартапам успішно залучити інвестиції, уникнути поширених помилок і побудувати стійкі бізнеси..

Книги та ідеї Brad Feld тісно пов'язані з аналізом фінансових ризиків у стартапах. Вони допомагають підприємцям оцінити фінансові ризики, пов'язані із залученням капіталу, уникнути втрати контролю над компанією та побудувати довгострокові стосунки з інвесторами, що сприяють стабільності бізнесу. Детальні рекомендації щодо переговорів, юридичних аспектів і фінансового планування роблять його роботи незамінними для будь-якого стартапу, що прагне досягти успіху в сучасному бізнес-середовищі.

2 МЕТОДИКА АНАЛІЗУ ТА ПРОГНАЗУВАННЯ РИЗИКІВ НА ОСНОВІ МАШИНОГО НАВЧАННЯ

2.1 Аналіз існуючих методик прогнозування ризиків на основі машинного навчання

2.1.1 Класичні методи аналізу та прогнозування ризиків

Регресійний аналіз є одним із фундаментальних статистичних методів, який широко застосовується для моделювання та кількісного оцінювання зв'язків між змінними. Його сутність полягає у вивченні залежності між однією залежною змінною (результатом, залежним показником) та набором незалежних змінних (факторів, предикторів), які впливають на цей результат. У контексті аналізу та прогнозування ризиків регресійний аналіз дозволяє чітко ідентифікувати ключові фактори, визначити ступінь їхнього впливу на ймовірність чи інтенсивність ризикових подій, а також розробити ефективні стратегії зниження або оптимізації ризиків. Завдяки регресійному аналізу управлінці, аналітики та дослідники отримують потужний інструмент для прийняття обґрунтованих рішень, що базуються не лише на інтуїції чи якісних оцінках, а й на кількісному вимірі впливу факторів, які стоять за тією чи іншою подією або результатом.

Основна формула лінійної регресії:

$$y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_n x_n + \varepsilon, \quad (2.1)$$

де y — залежна змінна (наприклад, рівень ризику, обсяг втрат, фінансовий показник, продуктивність чи інший цікавий індикатор);

x_i — незалежні змінні (фактори ризику, контрольні змінні, предиктори), які можуть бути кількісними (наприклад, вік, дохід, ціна) або якісними (категоріями, що кодуються за допомогою фіктивних змінних);

β_0 — вільний член моделі (інтерсепт), який відображає очікуване значення за у умови, що всі $x_i = 0$. Інтерсепт фактично задає відправну точку моделі;

β_i — коефіцієнти регресії, що кількісно відображають вплив кожної незалежної змінної x_i на залежну змінну y , за умови фіксованості інших факторів. Кожний коефіцієнт можна інтерпретувати як зміну середнього значення y при збільшенні x_i на одну одиницю;

ϵ — випадкова похибка, яка репрезентує невраховані фактори, шум у даних, помилки вимірювання та інші неконтрольовані аспекти. Вона вважається випадковою змінною із середнім значенням нуль та певною (часто незмінною) дисперсією.

Для коректної інтерпретації результатів необхідно забезпечити дотримання припущень регресійного аналізу. Серед основних припущень — лінійність зв'язку між залежною змінною та незалежними, незалежність похибок, відсутність автокореляції залишків, постійна дисперсія похибок (гомоскедастичність), нормальний розподіл похибок і відсутність мультиколінеарності серед предикторів. Порушення цих припущень може призвести до неадекватних оцінок параметрів, неправильних висновків та неточних прогнозів.

Для оцінки якості моделі та перевірки її адекватності використовують різноманітні статистичні критерії та показники. Коефіцієнт детермінації R^2 вказує, яку частку варіації залежної змінної пояснює побудована модель. Чим ближче R^2 до 1, тим краще модель підходить до даних. Однак високий R^2 сам по собі не гарантує коректності моделі, оскільки можна підвищити R^2 за рахунок додавання великої кількості предикторів, навіть якщо вони мало інформативні. Для порівняння моделей із різною кількістю предикторів застосовують інформаційні критерії, наприклад, критерій Акаїке або Байєсовський інформаційний критерій, які враховують також складність моделі.

Приклади компонентів регресійного аналізу

Компонент	Опис
Залежна змінна (Y)	Цільовий показник, який потрібно пояснити або спрогнозувати
Незалежні змінні (X)	Фактори або предиктори, що впливають на залежну змінну
Коефіцієнт (β)	Параметри моделі, що кількісно відображають вплив кожного фактора на результат
Інтерсепт (β_0)	Значення Y при X=0, відправна точка моделі
Похибка (ϵ)	Випадкова складова, яка враховує невраховані фактори та шум у даних
R^2	Коефіцієнт детермінації, частка поясненої варіації Y

Одним із важливих питань у регресійному аналізі є мультиколінеарність, коли незалежні змінні сильно корелюють між собою. Це ускладнює інтерпретацію коефіцієнтів і може призвести до нестабільності оцінок. Для виявлення мультиколінеарності використовують матрицю кореляцій, а також фактор інфляції дисперсії (VIF). Якщо мультиколінеарність надто висока, можна виключити одну зі змінних, застосувати методи регуляризації (Lasso, Ridge) або використати метод головних компонент для зменшення розмірності.

Якщо залишки моделі не відповідають припущенням (наприклад, виявлено

гетероскедастичність або автокореляцію), можуть застосовуватися робастні оцінки стандартних помилок, трансформації змінних або моделі зваженого найменшого квадрата. У разі виявлення викидів (аномальних спостережень), що надмірно впливають на оцінки коефіцієнтів, можна застосувати робастні методи регресії, які менш чутливі до окремих аномальних точок.

Окрім простих лінійних моделей, регресійний аналіз може бути розширений до складніших моделей, наприклад, поліноміальних регресій (для моделювання нелінійних зв'язків), регресійних моделей із взаємодією між змінними (коли ефект одного фактора залежить від рівня іншого фактора) та панельних регресій (для аналізу даних, що мають як часовий, так і крос-секційний вимір). Ці розширення дозволяють гнучкіше моделювати реальні складні системи, у яких відносини між змінними можуть бути далекі від простих лінійних патернів.

Застосування регресійного аналізу у сфері управління ризиками є надзвичайно широким. Його можна використовувати для оцінки ризиків кредитування (моделювання ймовірності дефолту), визначення факторів, що впливають на зміни курсів валют чи цін на сировину (ринкові ризики), оцінки впливу змін регуляторного середовища, моделювання ризиків у страхуванні (прогноз частоти страхових випадків та збитковості), виявлення факторів, що підвищують ризик нещасних випадків у виробничій сфері, прогнозування втрат у надзвичайних ситуаціях та багато іншого.

Завдяки можливості кількісно оцінити вплив різних чинників, регресійний аналіз стає основою для побудови економіко-математичних моделей ризиків, що допомагають стратегічно планувати діяльність підприємств, оптимізувати портфелі інвестицій, розраховувати страхові тарифи, приймати рішення щодо хеджування та диверсифікації.

В сучасному інформаційному середовищі, де обсяги даних постійно зростають, регресійний аналіз залишається одним із ключових інструментів. Його можна застосовувати до "великих даних", використовуючи паралельні обчислювальні потужності та оптимізовані алгоритми, що дозволяє швидко і ефективно оцінювати складні моделі з величезною кількістю спостережень та

змінних. Водночас важливо пам'ятати про інтерпретацію результатів: хоча регресійний аналіз може показувати, що певні фактори сильно пов'язані із залежною змінною, він не обов'язково встановлює причинно-наслідкові зв'язки. Для підтвердження причинності необхідні додаткові дослідження, експерименти або застосування спеціальних методів (наприклад, інструментальних змінних).

Метод найменших квадратів (МНК) є класичним і фундаментальним підходом до оцінювання параметрів регресійних моделей. Він базується на ідеї знаходження таких значень коефіцієнтів β_i , за яких сума квадратів різниць між фактичними значеннями залежної змінної та значеннями, передбаченими моделлю, стає мінімальною. Це дозволяє побудувати лінійну модель, яка найкраще відповідає наявним даним, у сенсі мінімуму середньоквадратичної помилки. МНК забезпечує отримання незсунених, ефективних та найкращих лінійних незсунених оцінок параметрів, коли виконуються класичні припущення регресійного аналізу (лінійність, нормальність похибок, гомоскедастичність, відсутність автокореляції та мультиколінеарності).

Математично задача оцінки коефіцієнтів β_i зводиться до мінімізації цільової функції:

$$\min_{\beta} \sum_{i=1}^N (y_i - \hat{y}_i)^2 \rightarrow 0, \quad (2.2)$$

де: y_i — фактичні значення залежної змінної для спостереження i ,

$\hat{y}_i$ — прогнозовані моделлю значення, які можна обчислити за формулою:

$$\hat{y}_i = \beta_0 + \beta_1 x_{i1} + \dots + \beta_n x_{in}. \quad (2.3)$$

Оцінки коефіцієнтів, отримані методом найменших квадратів, мінімізують суму квадратів відхилень $(y_i - \hat{y}_i)$, що забезпечує найбільш близьке підлаштування моделі під емпіричні дані у лінійному сенсі. Завдяки цьому ми отримуємо параметри, які роблять модель оптимальною для наявних спостережень, припускаючи, що змінні пов'язані лінійною залежністю.

Регресійний аналіз дозволяє визначити не лише величину впливу кожного фактора, але й статистичну значущість цього впливу. Оцінки коефіцієнтів β_i супроводжуються інформацією про їх стандартні помилки, що дає змогу проводити статистичні тести (наприклад, t-тести) для перевірки гіпотез про значущість параметрів. Використання р-значень дозволяє швидко оцінити, наскільки вплив певної змінної є істотним. Якщо р-значення є малим (наприклад, менше 0.05), це свідчить про статистичну значущість впливу змінної на залежну змінну. Такий підхід дозволяє аналітику або досліднику відсіяти незначущі змінні та спростити модель, зберігаючи при цьому лише фактори, які дійсно впливають на результат.

Окрім класичної лінійної регресії, широко застосовується логістична регресія, яка використовується у випадках, коли залежна змінна є категоріальною, а часто — бінарною (наприклад, настання події проти її ненастання). Замість прямого моделювання у, логістична регресія оцінює ймовірність настання події р, перетворюючи її за допомогою логіт-функції:

$$\ln\left(\frac{p}{1-p}\right) = \beta_0 + \beta_1 x_1 + \dots + \beta_n x_n. \quad (2.4)$$

Таке перетворення дозволяє моделювати ймовірність події в межах від 0 до 1, що є важливим для багатьох практичних завдань. Логістична регресія широко використовується у фінансах для оцінки ймовірності дефолту позичальників, у медицині для прогнозування наявності захворювання на основі симптомів та діагностичних показників, у маркетингу для передбачення ймовірності купівлі клієнтом певного товару чи послуги.

Однією з ключових переваг логістичної регресії є можливість працювати з факторними ознаками різної природи — як кількісними, так і якісними (категоріальними). Для категоріальних змінних використовують фіктивні змінні, які кодують категорії у вигляді 0 та 1. Крім того, результат логістичної моделі інтерпретується як зміна логіт-відношення шансів, що забезпечує більш зручну інтерпретацію у порівнянні з нелінійними показниками.

Логістична регресія також дозволяє оцінювати якість моделі, перевіряючи, наскільки добре вона класифікує або прогнозує настання події. Для цього використовують різні показники, такі як точність класифікації, повнота, специфічність, коефіцієнти-розходження, інформаційні критерії та інші. Аналітик може порівнювати різні моделі, змінюючи набір предикторів або перетворюючи дані, з метою поліпшення точності та стабільності прогнозів.

Таблиця 2.2

Порівняння лінійної та логістичної регресії

Характеристика	Лінійна регресія	Логістична регресія
Залежна змінна	Кількісна, неперервна	Двійкова, категоріальна (0 або 1)
Мета	Моделювання середнього значення Y при заданих X	Моделювання ймовірності події (p)
Функціональна форма	Лінійна комбінація X	Логіт-трансформація ймовірності p
Коефіцієнти (β)	Інтерпретуються як зміна Y при зміні X на одну одиницю	Інтерпретуються як зміна логіт-відношення шансів при зміні X на одну одиницю

Продовження таблиці 2.2

Порівняння лінійної та логістичної регресії

Характеристика	Лінійна регресія	Логістична регресія
Область значень	Y може бути будь-	p лежить в інтервалі

виходу	яким дійсним числом	(0; 1)
Застосування	Прогнозування кількісних показників (прибуток, обсяг виробництва)	Прогнозування подій: дефолт чи ні, наявність захворювання, вибір продукту, поведінка клієнта

Як при лінійному, так і при логістичному регресійному аналізі важлива ретельна перевірка моделі, включаючи діагностику залишків, оцінку мультиколінеарності, гетероскедастичності (у лінійній регресії), а також правильність вибору предикторів. Доцільно використовувати інформаційні критерії для порівняння моделей, тестувати різні комбінації змінних, а за потреби застосовувати методи регуляризації для боротьби з перенавчанням та покращення узагальнюючої здатності.

Одним Завдяки гнучкості, лінійна та логістична регресії можуть бути доповнені або поєднані з іншими статистичними і машинно-навчальними методами. У реальних завданнях можна використовувати трансформації змінних, розширені моделі (наприклад, додавання нелінійних членів, взаємодій між змінними) та застосовувати багатокрокові процедури вибору змінних. Також існують узагальнені лінійні моделі (GLM), які включають лінійну регресію, логістичну регресію, пуассонівську регресію та інші моделі, що дозволяють працювати з різними типами залежних змінних.

Для оцінки якості побудованої регресійної моделі використовується низка статистичних показників та критеріїв, які допомагають зрозуміти, наскільки добре модель відповідає наявним даним, чи можна їй довіряти при прогнозуванні та інтерпретації результатів, а також які аспекти слід покращити для підвищення її точності та надійності. Кожен із показників має своє призначення і доповнює загальну картину, дозволяючи досліднику приймати виважені рішення щодо подальших кроків в аналізі.

Одним із найпопулярніших показників є коефіцієнт детермінації (R^2). Він

вказує на частку варіації залежної змінної, яка може бути пояснена побудованою моделлю. Якщо коефіцієнт детермінації близький до одиниці, це свідчить про те, що модель добре підходить до даних. Проте занадто високий R^2 може свідчити про перенавчання моделі.

Таблиця 2.3

Інтерпретація значень R^2

Значення R^2	Інтерпретація
Близько до 0	Модель майже не пояснює варіацію залежної змінної
Помірне	Модель пояснює певну частку варіації, є інформативною, але не досконалою
Близько до 1	Модель дуже добре підходить до даних, але є ризик перенавчання

Стандартизована помилка регресії дозволяє оцінити середній розмір відхилення прогнозованих значень від фактичних, що допомагає зрозуміти точність та надійність прогнозів моделі. Цей показник схожий на середньоквадратичну помилку та часто використовується для порівняння моделей із різною шкалою залежної змінної.

Тест Фішера застосовується для перевірки загальної значущості моделі. Якщо результат тесту свідчить про статистично значущу модель, це означає, що набір предикторів має суттєвий спільний вплив на залежну змінну. Якщо ж тест Фішера не підтверджує значущості, можливо, модель не суттєво покращує прогноз у порівнянні з найпростішою моделлю, яка використовує лише середнє значення залежної змінної.

Аналіз залишків є невід'ємною частиною діагностики моделі. Залишки — це різниці між фактичними та прогнозованими значеннями залежної змінної. Якщо в них спостерігаються закономірності, це може свідчити про порушення припущень регресійного аналізу.

Таблиця 2.4

Типові проблеми, виявлені аналізом залишків

Проблема	Ознаки у залишках	Можливі рішення
Гетероскедастичність	Варіація залишків зростає чи зменшується з часом	Робастні оцінки, трансформація змінних
Автокореляція залишків	Залишки утворюють патерн (наприклад, хвилеподібний)	Використання моделей із серійною кореляцією
Нелінійність	Систематичні патерни в залишках	Додавання поліноміальних членів або сплайнів

Важливо перевіряти низку припущень регресійного аналізу. Якщо ці припущення порушені, застосовувані статистичні висновки можуть бути некоректними, а оцінки параметрів — ненадійними. Зокрема, слід переконатися у лінійності зв'язку, нормальності розподілу залишків, гомоскедастичності, відсутності автокореляції та мультиколінеарності.

Таблиця 2.6

Основні припущення та їх наслідки при порушенні

Припущення	Опис	Наслідки порушення
Лінійність	Залежність Y від X є лінійною	Неправильна специфікація моделі, систематичні патерни в залишках

Нормальність залишків	Залишки мають нормальний розподіл	Стандартні тести статистичної значущості можуть бути неточними
Гомоскедастичність	Дисперсія похибок є сталою	Неадекватні стандартні помилки, ненадійні інтервали довіри

Таблиця 2.7

Основні припущення та їх наслідки при порушенні

Припущення	Опис	Наслідки порушення
Відсутність автокореляції	Похибки є незалежними	Систематичні помилки у прогнозах, заниження чи завищення точності
Відсутність мультиколінеарності	Незалежні змінні не сильно корелюють між собою	Нестабільність оцінок, проблеми з інтерпретацією коефіцієнтів

Додаткові методи стають в пригоді, коли класичний лінійний підхід не дає задовільних результатів або коли модель потребує уточнень. Нелінійні регресії використовуються, якщо залежність між змінними не може бути адекватно описана лінійною функцією. Панельні регресії потрібні тоді, коли дані мають часовий та крос-секційний вимір, що дає змогу враховувати індивідуальні особливості об'єктів та їхню динаміку. Регуляризаційні методи (Lasso, Ridge) допомагають при наявності великої кількості змінних та проблемі мультиколінеарності або перенавчання. Вони додають штраф за складність моделі, стимулюючи виключення несуттєвих змінних.

Таблиця 2.8

Додаткові методи модифікації та вдосконалення моделі

Метод	Ситуації застосування	Очікуваний ефект
Нелінійні регресії	Коли зв'язок між змінними не є лінійним	Краще моделювання складних залежностей
Панельні регресії	Дані з часовим та крос-секційним виміром	Врахування індивідуальних особливостей та динаміки
Регуляризаційні методи	Велика кількість змінних, мультиколінеарність	Запобігання перенавчанню, спрощення моделі

Байєсові моделі ґрунтуються на теоремі Байєса та пропонують унікальний статистичний та концептуальний підхід до інтерпретації ймовірності, моделювання невизначеності та оновлення знань.

Відмінність від класичних фреквентистських методів полягає в тому, що класична статистика розглядає параметри як фіксовані, але невідомі величини, а ймовірність подій — як граничну частоту в нескінченній послідовності

експериментів. Натомість байєсовий підхід інтерпретує ймовірність як суб'єктивну оцінку ступеня впевненості в гіпотезі чи параметрі.

Це означає, що параметри моделі розглядаються як випадкові змінні з апіорними розподілами, які відображають наші знання або упередження до отримання нових даних. Коли з'являється нова інформація, ці апіорні розподіли оновлюються, перетворюючись на апостеріорні розподіли, які враховують як початкові уявлення, так і емпіричну інформацію, здобуту з даних..

Такий підхід особливо цінний в умовах невизначеності або браку інформації. Наприклад, у ситуаціях, коли дані є дорогими для отримання, коли ми маємо лише невеликий вибіркового обсяг чи коли система є надзвичайно складною, байєсові моделі дозволяють інтегрувати попередні експертні оцінки, історичні дані чи теоретичні уявлення. Це допомагає отримати більш адекватні оцінки параметрів навіть за нестачі інформації. Завдяки тому, що ймовірність інтерпретується суб'єктивно, дослідник має гнучкість у формуванні апіорів, обираючи їх так, щоб вони відображали наявні знання чи упередження, а потім оновлювати ці апіорі при надходженні нових даних.

Основоположним інструментом у цьому процесі є теорема Байєса, яка пропонує формальний механізм оновлення ймовірностей:

$$P\left(\frac{H}{D}\right) = \frac{P\left(\frac{D}{H}\right) \cdot P(H)}{P(D)}, \quad (2.5)$$

де: H — гіпотеза, яку ми оцінюємо;

D — нові дані, які ми отримуємо з реальних спостережень, експериментів чи вимірювань;

$P(H)$ — апіорна ймовірність гіпотези H , яка відображає нашу впевненість у цій гіпотезі до врахування нових даних;

$P(D|H)$ — правдоподібність, тобто ймовірність спостереження даних D , за умови, що гіпотеза H істинна. Цей компонент відображає, наскільки добре гіпотеза пояснює отримані дані;

$P(D)$ — загальна ймовірність спостереження даних, що виступає нормувальним коефіцієнтом. Її можна отримати інтегруючи або сумуючи за всіма можливими гіпотезами;

$P(D)$ гарантує, що сума апостеріорних ймовірностей всіх гіпотез дорівнює.

Таблиця 2.9

Переваги та обмеження моделей Байєса

Аспект	Переваги	Обмеження
Врахування попередніх знань	Інтегрує історичні/експертні дані	Чутливість до вибору апріорів
Гнучкість	Модель складних систем	Високі обчислювальні затрати
Повний опис невизначеності	Повний розподіл параметрів	Складні МСМС-метод

Таблиця 2.10

Переваги та обмеження моделей Байєса

Оновлюваність	Динамічне оновлення з новими даними	Потрібен досвід у байєсовій статистиці
Широке застосування	Багато сфер використання	Потрібне спеціалізоване ПЗ

Аналіз дерев рішень — це методологія, що надає можливість структурно та послідовно оцінювати складні рішення в умовах невизначеності, коли перед дослідником або особою, яка приймає рішення, постає низка можливих варіантів дій, а результат кожного з них залежить не лише від обраної стратегії, але й від випадкових подій із певними ймовірностями.

Такий підхід дозволяє чітко уявити послідовність рішень і подій, буде інтуїтивну візуальну модель у вигляді дерева, де вузли рішень відображають точки, в яких треба зробити вибір, а вузли випадковостей — ситуації, коли результат визначається випадком. Кожна гілка дерева відображає окремий сценарій чи варіант розвитку подій, і, рухаючись цими гілками, можна оцінити можливі наслідки у вигляді вигод, витрат або прибутків.

Листя дерева, що розташовуються на кінцях гілок, показують кінцеві результати, які можуть бути як сприятливими, так і несприятливими, залежать від логіки подій, стратегії дій та ймовірностей різних непередбачених ситуацій. Таким чином, у межах байєсового підходу параметри розглядаються не як фіксовані числа, а як випадкові змінні із власними розподілами ймовірностей. Це відкриває широкі можливості для аналізу невизначеності, адаптації до нової інформації та моделювання складних ситуацій, де класичні методи можуть виявитися недостатньо гнучкими. Байєсові моделі забезпечують багатий інструментарій для інкорпорації апріорних знань, обробки складних структур даних та створення інформативних висновків, що можуть підвищити точність та корисність статистичних аналізів у багатьох галузях діяльності.

Таблиця 2.11

Основні елементи дерева рішень

Елемент	Опис
Вузли рішень	Точки, де ухвалюється конкретне рішення
Вузли випадковостей	Точки, де результат визначається випадковою подією з певною ймовірністю
Гілки	Лінії, що представляють можливі варіанти дій або результати подій
Листя (кінцеві вузли)	Кінцеві результати з відповідними вигодами або витратами

Однією з ключових переваг застосування дерев рішень є їх наочність.

Навіть досить складні проблеми можна розкласти на логічні кроки, а потім представити у вигляді зрозумілого графічного образу. Ця візуалізація допомагає швидко охопити сутність проблеми та пояснити її іншим, обґрунтувавши вибір певної стратегії. Крім того, дерева рішень надають можливість урахування ймовірностей та очікуваних значень результатів, переходячи від інтуїтивної оцінки сценаріїв до їх кількісного аналізу.

Ключовим інструментом для оцінки сценаріїв у вузлах випадковостей є обчислення очікуваного значення (EV), яке дозволяє інтегрувати інформацію про ймовірності можливих результатів та величини вигод чи втрат.

Формула для обчислення очікуваного значення у вузлі:

$$EV = \sum_{i=1}^k p_i \cdot V_i, \quad (2.6)$$

де: EV — очікуване значення у вузлі випадковостей.

p_i — ймовірність настання i -го результату, пов'язаного з певною подією чи сценарієм;

V_i — значення (вигода, витрата, прибуток, корисність або будь-який інший оціночний показник), яке відповідає i -му результату;

За цією формулою можна розрахувати EV для будь-якого вузла випадковостей, після чого ця величина може бути використана для "підйому" по дереву вгору, до вузлів рішень. На вузлах рішень, маючи обчислені EV для кожної гілки, можна обрати ту гілку, яка дає максимальне очікуване значення, тобто буде оптимальним вибором з погляду раціонального максимізатора вигоди чи мінімізатора ризику.

Типові завдання та можливості аналізу дерев рішень

Завдання	Можливості, що надає аналіз дерев рішень
Оцінка інвестиційних проектів	Порівняння стратегій, розрахунок очікуваних прибутків чи збитків
Вихід на новий ринок	Врахування конкурентних реакцій, економічних умов, оцінка ймовірностей різних сценаріїв
Планування маркетингових кампаній	Визначення оптимального каналу збуту, оцінка очікуваного впливу на продажі
Медичні рішення	Порівняння лікувальних стратегій, прогнозування шансів успіху чи ускладнень

Проте, із застосуванням дерев рішень виникає низка потенційних проблем. По-перше, якщо проблема дуже складна, дерево може стати надто розгалуженим та громіздким, а його побудова й аналіз стануть трудомісткими. По-друге, точність результатів значною мірою залежить від адекватності введених ймовірностей та оцінок величин вигод чи втрат. Якщо ці оцінки надто приблизні або хибні, модель може призвести до неправильних висновків. По-третє, стандартний аналіз дерев рішень часто передбачає незалежність подій чи спрощені припущення, тому якщо реальна система має складні кореляції між факторами, можливо, доведеться розглянути додаткові або альтернативні методи моделювання, такі як мережі Байєса чи імітаційне моделювання.

Аналіз чутливості є потужним інструментом у процесі оцінки моделей, який дозволяє зрозуміти, як зміни вхідних параметрів впливають на результати моделі. Це важливо для визначення найбільш критичних факторів, які мають найбільший

вплив на кінцевий результат, а також для оптимізації управлінських рішень, спрямованих на зменшення ризиків та покращення продуктивності системи. Аналіз чутливості допомагає виявити, які саме параметри слід контролювати найбільш ретельно, оскільки вони суттєво впливають на результативність моделі.

Сценарний аналіз є одним із найефективніших інструментів у процесі оцінки ризиків та прийняття стратегічних рішень. Він дозволяє моделювати різні можливі майбутні події та оцінювати їхній вплив на результати проєкту чи організації. Це особливо важливо у середовищах з високою невизначеністю, де традиційні методи прогнозування можуть бути недостатньо точними або адаптивними. Сценарний аналіз допомагає підготуватися до різних варіантів розвитку подій, забезпечуючи основу для розробки гнучких та адаптивних стратегій управління ризиками.

Аналіз базується на концепції створення кількох можливих сценаріїв розвитку подій, кожен з яких відображає певний набір умов та факторів. Ці сценарії можуть бути оптимістичними, песимістичними або базовими, що дозволяє охопити широкий спектр можливих майбутніх станів. Основна мета — зрозуміти, як різні комбінації факторів впливають на результати, та визначити найбільш критичні аспекти, які потребують уваги та управління.

Сценарний аналіз складається з кількох ключових етапів, кожен з яких виконує специфічну роль у формуванні повноцінного аналізу ризиків:

- Визначення ключових факторів невизначеності
- Розробка логічно послідовних сценаріїв
- Оцінка впливу кожного сценарію на результати
- Вибір методів та інструментів для сценарного аналізу
- Розробка стратегій управління ризиками

Приклад ключових факторів ризиків невизначеності

Ключові фактори невизначеності	Опис
Темпи зростання ринку	Швидкість зростання попиту на продукти чи послуги стартапу
Залучення нових клієнтів	Способи та ефективність залучення нових користувачів
Рівень інвестицій	Обсяг та стабільність фінансування стартапу
Технологічні зміни	Інновації та технічні покращення в продуктах стартапу
Конкурентне середовище	Кількість та сила конкурентів на ринку

Ці фактори слугують основою для подальшого розроблення сценаріїв, які відображають різні можливі майбутні стани ринку та внутрішні умови стартапу.

Розробка логічно послідовних сценаріїв включає створення кількох сценаріїв, кожен з яких представляє певний набір умов та подій. Вони повинні бути логічно послідовними та реалістичними, відображаючи різні можливі майбутні стани.

Оцінка впливу кожного сценарію на результати проводиться аналіз того, як кожен сценарій впливатиме на ключові результати проєкту. Це включає використання математичних моделей, фінансових прогнозів та інших інструментів для оцінки можливих наслідків.

Для проведення сценарного аналізу використовуються різноманітні методи, які допомагають структурувати процес, забезпечувати точність аналізу та оптимізувати робочі процеси.

Таблиця 2.14

Основні методи сценарного аналізу та їх характеристики

Метод	Опис	Переваги	Обмеження
Монте-Карло	Генерація великої кількості випадкових сценаріїв	Глобальна оцінка чутливості	Висока обчислювальна витрата
Системна динаміка	Моделювання складних взаємодій між компонентами системи	Враховує динаміку змін	Складність моделювання
Дерева рішень	Візуалізація різних сценаріїв та їхніх наслідків	Легко зрозумілий, візуально привабливий	Обмежена точність при складних взаємодіях
Динамічні моделі	Аналіз змін параметрів протягом часу	Дозволяє прогнозувати довгострокові наслідки	Вимагає складних математичних моделей

Сценарний аналіз може бути ефективно поєднаний з іншими методами аналізу ризиків, такими як аналіз чутливості та байєсові методи, для отримання більш повної та глибокої картини ризиків.

Сценарний аналіз є невід'ємною частиною стратегічного планування, оскільки він дозволяє організаціям прогнозувати можливі майбутні події та підготуватися до них заздалегідь. Це допомагає мінімізувати ризики, оптимізувати ресурси та забезпечити стійкість організації у разі непередбачених змін.

Для кількісної оцінки ефективності сценарного аналізу можна використовувати метрики, що вимірюють, наскільки добре сценарії покривають можливі майбутні події та як точно вони прогнозують результати.

Формула (2.7) для оцінки покриття сценаріїв:

$$C = \frac{N_{covered}}{N_{total}} \cdot 100\%, \quad (2.7)$$

де C — покриття сценаріїв;

$N_{covered}$ — кількість охоплених сценаріїв;

N_{total} — загальна кількість можливих сценаріїв.

Аналіз часових рядів є важливим інструментом у статистиці та економіці, який використовується для моделювання та прогнозування змін різних показників у часі. Цей метод дозволяє виявити та зрозуміти структурні компоненти даних, такі як тренди, сезонність, циклічність та випадковість. Розуміння цих компонентів є ключовим для прийняття обґрунтованих рішень у бізнесі, фінансах, охороні здоров'я та інших сферах, де важливо прогнозувати майбутні події та тенденції.

Часові ряди складаються з кількох ключових компонентів, які визначають їхню поведінку та дозволяють ефективно аналізувати та прогнозувати майбутні значення.

Тренд відображає загальний напрямок зміни показника протягом тривалого періоду часу, може бути як зростаючим, так і спадним. Наприклад, постійне зростання продажів компанії протягом кількох років свідчить про позитивний тренд.

Сезонність характеризується повторюваними коливаннями через певні періоди часу, такі як місяці, квартали або роки. Наприклад, підвищений попит на зимовий одяг перед зимою є типовим прикладом сезонності.

Циклічність включає коливання з більшою тривалістю, ніж сезонність, часто пов'язані з економічними циклами. Наприклад, економічні підйоми та спади протягом десятиліть відображають циклічність у даних.

Випадковість представляє невизначені коливання, які не можна пояснити іншими компонентами. Це можуть бути несподівані події, такі як природні катастрофи або технічні збої, що впливають на значення ряду без очевидної причинної зв'язку.

Для аналізу та прогнозування часових рядів використовуються різні

математичні моделі, кожна з яких має свої переваги та обмеження.

Авторегресійна модель (AR) передбачає, що поточне значення часової серії залежить від її попередніх значень. Ця модель використовується для опису залежності між спостереженнями та визначення автокореляції у даних.

Формула (2.8) AR(p):

$$Y_t = c + \phi_1 Y_{t-1} + \phi_2 Y_{t-2} + \dots + \phi_p Y_{t-p} + \epsilon_t, \quad (2.8)$$

де: Y_t — поточне значення серії;

c — константа;

$\phi_1, \phi_2, \dots, \phi_p$ коефіцієнти авторегресії;

ϵ_t — випадкова похибка.

Модель ковзного середнього (MA) передбачає, що поточне значення серії залежить від попередніх випадкових похибок. Ця модель використовується для корекції систематичних помилок у прогнозах.

Формула (2.9) MA(q):

$$Y_t = c + \epsilon_t + \theta_1 \epsilon_{t-1} + \theta_2 \epsilon_{t-2} + \dots + \theta_q \epsilon_{t-q}, \quad (2.9)$$

де: Y_t — поточне значення серії;

c — константа;

$\theta_1, \theta_2, \dots, \theta_q$ — коефіцієнти ковзного середнього;

ϵ_t — випадкова похибка.

Модель ARIMA поєднує авторегресійні (AR) та ковзні середні (MA) компоненти, а також інтегрування для роботи з нестационарними рядами. Ця модель є однією з найбільш популярних для аналізу часових рядів завдяки своїй гнучкості та здатності моделювати різні види даних.

Формула (2.10) ARIMA(p, d, q):

$$\Phi(B)(1 - B)^d Y_t = \theta(B)\epsilon_t, \quad (2.10)$$

де p — порядок авторегресії;

d — порядок диференціювання;

q — порядок ковзного середнього;

B — лаг-оператор.

Імітаційне моделювання є потужним інструментом аналізу та оптимізації складних систем у різних галузях, таких як логістика, виробництво, фінанси, охорона здоров'я та екологія. Воно дозволяє створити віртуальну модель реальної системи, експериментувати з нею, змінювати параметри та умови, а також прогнозувати поведінку системи в різних сценаріях. Це особливо корисно для систем, де аналітичні рішення важко отримати через високу складність або непередбачуваність взаємодій між компонентами.

Імітаційне моделювання охоплює кілька ключових аспектів, які визначають його ефективність та застосовність у різних сферах:

Створення моделі полягає у відтворенні ключових характеристик реальної системи у віртуальному середовищі. Це включає визначення основних компонентів системи, їх взаємодій, правил поведінки та динаміки зміни параметрів.

Експериментування з параметрами дозволяє змінювати різні аспекти моделі для дослідження їх впливу на загальну поведінку системи.

Аналіз результатів включає оцінку отриманих даних для прийняття обґрунтованих рішень щодо покращення системи.

Валідація моделі є критичним етапом, який забезпечує достовірність та надійність моделі.

Оптимізація та впровадження включає використання отриманих результатів для вдосконалення реальної системи.

Переваги та обмеження імітаційного моделювання

Категорія	Переваги	Обмеження
Аналіз	Аналіз складних систем	Вимогливість до обчислювальних ресурсів
Тестування	Можливість тестування різних сценаріїв	Залежність від точності моделі
Розуміння	Підвищення розуміння внутрішніх процесів системи	Може бути складним у реалізації
Рішення	Підтримка прийняття обґрунтованих рішень	Вимоги до якісних та повних даних

Хоча імітаційне моделювання зазвичай не базується на конкретних математичних формулах, воно використовує різні математичні та статистичні методи для опису поведінки системи. Ось кілька прикладів:

Формула (2.11) Час між подіями (Inter-arrival Time) у DES:

$$T_i = \frac{\ln(U)}{\lambda}, \quad (2.11)$$

де U — випадкове число з рівномірним розподілом $[0,1]$;

λ — інтенсивність потоку подій.

Формула (2.12) Генерація випадкових чисел для Монте-Карло симуляції:

$$X \sim N(\mu, \sigma^2), \quad (2.12)$$

де X — випадкова змінна з нормальним розподілом з параметрами середнього μ та дисперсії σ^2 .

Імітаційне моделювання є важливим інструментом сучасного управління та оптимізації систем у різних галузях. Воно дозволяє організаціям та проектам

краще розуміти внутрішні процеси, тестувати різні стратегії управління та приймати обґрунтовані рішення для підвищення ефективності та стійкості. Незважаючи на певні складнощі та обмеження, систематичне застосування імітаційного моделювання значно покращує процес планування та управління у складних та динамічних умовах.

2.2 Сучасні методи прогнозування ризиків з використанням машинного навчання

Сучасні методи прогнозування ризиків з використанням машинного навчання охоплюють широкий спектр підходів та алгоритмів, які дають змогу ефективно обробляти великі обсяги даних, враховувати складні взаємозв'язки між факторними ознаками, а також адаптуватися до динамічно змінюваного середовища. Серед найбільш успішних та поширених методів, що застосовуються для прогнозування ризиків у фінансах, страхуванні, промисловості, охороні здоров'я, енергетиці та інших сферах, можна виділити алгоритми опорних векторів (SVM), нейронні мережі та ансамблеві методи, такі як випадкові ліси (Random Forest) та градієнтний бустинг (Gradient Boosting).

Класифікація ризиків за допомогою алгоритмів підтримки векторів (Support Vector Machines, SVM) є одним із найбільш фундаментальних і водночас гнучких інструментів машинного навчання, здатних вирішувати складні задачі розмежування об'єктів за різними класами. SVM відзначаються тим, що вони не просто намагаються знайти будь-яку межу поділу між класами, а прагнуть знайти таку, яка максимізує відстань (маргін) до найближчих навчальних точок обох класів. Такий підхід забезпечує кращу узагальнюючу здатність моделі, оскільки вона менш схильна до перенавчання та може впевненіше прогнозувати клас нових, невідомих спостережень.

Основна ідея методу — знайти розділювальну гіперплощину, яка максимально відділяє точки одного класу від іншого. Якщо розглянути задачу лінійної класифікації, маємо навчальну вибірку:

$$\{(x_i, y_i)\}_{i=1}^N, \quad (2.13)$$

де $x_i \in R^2$ — вектор ознак i -го спостереження;

$y_i \in \{-1, 1\}$ — мітка класу, яка вказує, чи належить об'єкт до ризикової категорії;

V_i — значення (вигода, витрата, прибуток, корисність або будь-який інший оціночний показник), яке відповідає i -му результату.

Таблиця 2.16

Типи ядерних функцій SVM

Тип ядра	Приклад ядрової функції	Застосування
Лінійне	$k(x_i, x_j) = x_i \cdot x_j$	Для проблем із лінійними залежностями
Поліноміальне	$k(x_i, x_j) = (x_i \cdot x_j + c)^2$	Якщо залежність можна наблизити поліномом
RBF (Радіально-базисне)	$k(x_i, x_j) = \exp(-\gamma \ x_i - x_j\ ^2)$	Популярне для складних нелінійних залежностей
Сигмоїдне	$k(x_i, x_j) = \tan(ax_i \cdot x_j + c)$	Іноді нагадує поведінку нейронних мереж

SVM добре працює при високій розмірності простору ознак, коли класичні методи можуть страждати від "прокляття розмірності". Він також є досить стійким до викидів чи шуму, оскільки оптимізація спрямована на пошук максимальної відстані від гіперплощини до найближчих точок, що зменшує вплив окремих аномалій. Важливо також, що SVM забезпечує контроль над балансом між помилками навчання та узагальнюваністю завдяки використанню параметрів

штрафу за помилки класифікації (C) та налаштуванню параметрів ядерних функцій.

У практичних задачах прогнозування ризиків застосування SVM може дати досить високі показники точності класифікації. Наприклад, у фінансовій сфері можна використовувати SVM для ідентифікації клієнтів банку, які потенційно не повернуть кредит, чи виділення підозрілих транзакцій, що можуть вказувати на шахрайство. У технічних задачах можна застосувати SVM для класифікації станів обладнання, визначення деталей, що скоро можуть вийти з ладу, та запобігання аварійним ситуаціям.

Використання нейронних мереж для прогнозування складних сценаріїв є важливим кроком уперед у підвищенні точності, гнучкості та адаптивності моделей, що прогнозують ризики. Традиційні статистичні та машинно-навчальні підходи часто вимагають ручного вибору ознак або спеціальної обробки даних, щоб витягнути корисну інформацію. Натомість нейронні мережі, особливо глибокі архітектури глибинного навчання (Deep Learning), здатні самостійно навчатися релевантним ознакам без прямого втручання дослідника. Це особливо корисно, коли дані складні, багатовимірні, слабо структуровані або включають нетрадиційні формати інформації.

Сучасні дані, що відображають ризики, можуть бути дуже різноманітними. Наприклад, у сфері прогнозування фінансових ризиків важливо враховувати історичні часові ряди з показниками ринку, опрацьовувати текстові документи (контракти, звіти, новини), аналізувати зображення (фотографії стану обладнання, інфраструктури), а інколи навіть сигнали (аудіозаписи, сенсорні дані). Нейронні мережі пропонують механізми для інтеграції цих різноманітних джерел інформації в єдину модель.

Наприклад, згорткові нейронні мережі (CNN) ефективні для роботи із зображеннями та можуть бути адаптовані до обробки послідовних даних. Рекурентні архітектури (RNN, LSTM, GRU) розроблені для аналізу часових рядів та послідовностей, а трансформери дозволяють ефективно працювати з текстами та складними структурами взаємодій між ознаками, запам'ятовуючи

довгострокові залежності і при цьому бути більш обчислювально ефективними порівняно з класичними рекурентними мережами.

Базовим елементом є штучний нейрон, який здійснює лінійну комбінацію вхідних ознак та передає її через нелінійну активаційну функцію. Якщо розглянути один шар мережі, маємо:

$$h = f(W_x + b), \quad (2.14)$$

де x — вектор вхідних ознак;

W — матриця ваг;

b — вектор зсуву;

$f()$ — нелінійна функція активації.

Нейронна мережа зазвичай складається з кількох таких шарів, утворюючи глибинну архітектуру:

$$y = f^{(L)} \left(\dots f^{(2)} \left(f^{(1)}(x) \right) \dots \right), \quad (2.15)$$

де $f^{(l)}$ — відображення, що здійснюється L -им шаром.

Це дозволяє мережі послідовно витягувати дедалі складніші та високорівневі ознаки, починаючи від простих і переходячи до складних. Наприклад, у разі зображень перші шари можуть виявляти контури, наступні — більш складні фрагменти, а глибинні шари — цілі об'єкти або патерни, пов'язані з ризиком. Для часових рядів початкові шари можуть виділяти короткострокові шаблони, а глибинні шари виявляють довгострокові тренди чи залежності.

Основні типи архітектур нейронних мереж та їх застосування до прогнозування ризиків

Архітектура	Опис	Приклади застосування
CNN (Згорткові мережі)	Обробка зображень, двовимірних сигналів, екстракція просторових патернів	Аналіз фотографій обладнання для виявлення ознак зносу чи пошкоджень
RNN, LSTM, GRU	Робота з часовими рядами та послідовними даними	Прогнозування фінансових ринкових індикаторів, виявлення патернів у динамічних системах
Трансформери	Обробка текстів, складних послідовностей, увага до важливих частин даних	Аналіз текстових звітів, статей, новин для прогнозування подій, що впливають на ризики
Аутоенкодери	Стиснення даних, навчання латентних представлень	Виявлення аномалій у даних, пошук нетипових ситуацій із підвищеним ризиком

Використання нейронних мереж у прогнозуванні ризиків дозволяє враховувати тонкі патерни та приховані залежності. Наприклад, якщо йдеться про прогнозування ймовірності дефолту позичальника, нейронна мережа може одночасно обробляти історичні фінансові показники клієнта, його транзакційну активність, інформацію з кредитної історії, а також текстові замітки спеціалістів чи зображення, які відображають стан заставного майна. Мережа здатна навчитися визначати ознаки, які вказують на погіршення платоспроможності, навіть якщо ці патерни були б важкими для виявлення традиційними методами.

Сучасні нейронні мережі, зокрема глибокі архітектури (Deep Learning), відзначаються здатністю автоматично витягувати релевантні ознаки з сирих даних, що має особливе значення у випадках, коли дані є слабо структурованими

або коли важко заздалегідь визначити, які ознаки будуть впливати на ризик. Наприклад, якщо модель має враховувати історичні часові ряди, текстові дані (звіти, контракти), зображення (фотографії обладнання), звуки чи інші типи сигналів, нейронні мережі дозволяють інтегрувати ці різноманітні джерела інформації в єдину модель.

Різні типи архітектур, такі як згорткові нейронні мережі (CNN) для роботи із зображеннями та послідовними даними, рекурентні нейронні мережі (RNN, LSTM, GRU) для опрацювання часових рядів і послідовностей, а також трансформери для аналізу текстів чи складних взаємодій, роблять нейронні мережі універсальним інструментом. Завдяки цим можливостям використання нейронних мереж у прогнозуванні ризиків допомагає враховувати тонкі патерни та приховані залежності, які були б важкими або неможливими для виявлення традиційними методами.

Ансамблеві методи, такі як Random Forest та Gradient Boosting, також відіграють ключову роль у сучасних підходах до прогнозування ризиків. Ідея ансамблевих методів полягає у тому, що замість побудови однієї потужної, але можливо нестійкої моделі, краще створити цілий "ансамбль" відносно простих моделей (наприклад, дерев рішень), а потім об'єднати їхні прогнози. Random Forest будується з великої кількості незалежних дерев рішень, кожне з яких навчається на випадковій підмножині ознак і спостережень. Середнє або мода прогнозів усіх цих дерев дає більш стабільну та точну модель, оскільки випадковість допомагає уникати перенавчання і підвищує узагальнюючу здатність..

Важливо, що сучасні методи машинного навчання дозволяють не лише підвищувати точність прогнозів, але й робити процес моделювання більш гнучким і адаптивним. За потреби можна поєднувати різні методи: використання предикторів, згенерованих глибокою нейронною мережею, як вхідних змінних для SVM чи Gradient Boosting, застосування ансамблів з різних типів моделей, регуляризація для запобігання перенавчанню, крос-валідація для об'єктивної оцінки якості моделей, а також засоби Explainable AI (XAI) для інтерпретації

складних прогнозів. Це забезпечує можливість створення систем прогнозування ризиків, які можуть бути не лише точними, але й прозорими настільки, наскільки цього вимагають реальні умови прийняття рішень.

У практичних задачах прогнозування ризиків застосування SVM може дати досить високі показники точності класифікації. У технічних задачах можна застосувати SVM для класифікації станів обладнання, визначення деталей, що скоро можуть вийти з ладу, та запобігання аварійним ситуаціям.

3 РОЗРОБКА МЕТОДУ АНАЛІЗУ ТА ПРАГНАЗУВАННЯ РИЗИКІВ У СТАРТАП -ПРОЄКТАХ НА ОСНОВІ ТЕХНОЛОГІЇ МАШИНОГО

3.1 Опис розробки методу

Розроблений метод прогнозування ризиків у стартап-проектах покликаний стати надійним, гнучким та інформативним інструментом для аналізу складних, багатофакторних процесів, які визначають ймовірність досягнення успіху чи провалу нового підприємства. Стартапи діють у мінливому середовищі, де відсутність історично стабільних часових рядів, невизначеність ринкової кон'юнктури, обмеженість ресурсів та швидкі організаційні зміни ускладнюють побудову класичних статистичних моделей. У таких умовах використання сучасних машинно-навчальних технологій дозволяє перейти від статичного уявлення про ризик до динамічного, адаптивного підходу, в якому знання оновлюються разом із появою нової інформації, а методи автоматично підлаштовуються до особливостей даних.

Серед ключових компонентів методу можна виділити три рівні інтеграції. Перший рівень охоплює застосування байєсових принципів та базових статистичних моделей. На цьому етапі ми використовуємо апріорні знання, отримані з експертних оцінок, історичних даних про інші стартапи, теоретичних уявлень про галузеві ризики та економічні чинники, а також робимо перший крок у формуванні первинних прогнозів.

Байєсові моделі дозволяють поєднувати суб'єктивну впевненість із даними, динамічно оновлювати апріорну ймовірність гіпотез про ризик із появою нової інформації. Це створює основу для прийняття рішень за умов, коли даних мало, але є цінні експертні знання, або коли система є настільки новою, що класичні методи без апріорів втрачають точність і стабільність.

У практичних задачах прогнозування ризиків застосування SVM може дати

досить високі показники точності класифікації. Наприклад, у фінансовій сфері можна використовувати SVM для ідентифікації клієнтів банку, які потенційно не повернуть кредит, чи виділення підозрілих транзакцій, що можуть вказувати на шахрайство. У технічних задачах можна застосувати SVM для класифікації станів обладнання, визначення деталей, що скоро можуть вийти з ладу, та запобігання аварійним ситуаціям

Другий рівень охоплює розширення простоти лінійних моделей та ймовірнісного підходу за рахунок алгоритмів машинного навчання, таких як логістична регресія (із регуляризацією та вибором ознак), алгоритми підтримки векторів (SVM) та нейронні мережі різних архітектур. Логістична регресія може слугувати вихідною моделлю для основних оцінок ризику — наприклад, ймовірності дефолту або ймовірності невиконання KPI стартапом. Вона легко інтерпретується, що на початковому етапі важливо для розуміння взаємозв'язків між факторами. Проте логістична регресія лінійна, і коли система є складною, цього може бути недостатньо.

Саме тут у гру вступають SVM із використанням нелінійних ядер, що дає змогу моделювати складні межі між класами "ризик" і "неризик". Якщо дані багатовимірні, містять приховані закономірності, які нелінійні та важко виявити, SVM із ядерними функціями дозволяє вищовимірне відображення ознак і лінійну сепарацію вже у цьому розширеному просторі. Це підвищує точність класифікації, дозволяючи тонко налаштувати розділову гіперплощину.

Далі нейронні мережі, зокрема глибокі, вносять додатковий рівень гнучкості. Завдяки їх здатності працювати з неструктурованими даними (текстами, зображеннями, сигналами) та великими наборами особливостей, модель стає здатною враховувати не лише класичні числові метрики, а й контекстні, динамічні та багаторівневі патерни, що впливають на ризик. Наприклад, якщо стартап працює у сфері штучного інтелекту, аналіз технологічних тенденцій у наукових публікаціях чи патентах можна здійснити через текстові дані.

Нейронна мережа (наприклад, трансформер) допоможе витягти корисну

інформацію з цих текстів та інтегрувати її в загальний прогноз ризику, поєднуючи з іншими ознаками, як-от показники швидкості росту команди, структури витрат і поведінкових індикаторів.

Третій рівень інтеграції — це застосування ансамблевих методів. Замість покладатися на одну модель, яка, можливо, занадто чутлива до шуму у даних, ми створюємо ансамбль різних моделей або варіацій однієї моделі. Наприклад, Random Forest утворюється з багатьох дерев рішень, кожне з яких бачить тільки випадковий піднабір ознак і спостережень. Такий "стохастичний" підхід мінімізує ризик того, що якийсь один набір факторів викривить прогноз. Усереднення або голосування прогнозів усіх цих дерев дає стабільний та узагальнюючий результат. Gradient Boosting покращує точність прогнозу, поступово додаючи нові моделі, кожна з яких фокусується на помилках попереднього ансамблю. Зрештою, отримуємо дуже гнучку та точну модель, що витягує максимум інформації з наявних даних.

Таким чином, остаточна методологія — це багаторівневий, інтегрований підхід, де використання різних методів здійснюється послідовно чи паралельно, залежно від характеристик задачі та структури даних. Наприклад, на початку можна застосувати процедури попередньої обробки (видалення пропусків, нормалізація, зменшення розмірності), потім використати логістичну регресію для початкового розуміння, далі SVM — для покращення нелінійної класифікації, нейронні мережі — для вилучення складних ознак із різномірних джерел, і, нарешті, ансамблеві методи (Random Forest, Gradient Boosting) — для підвищення стабільності та точності кінцевого прогнозу. Байєсові принципи можуть бути вбудовані на будь-якому етапі, дозволяючи оновлювати апріорні розподіли при надходженні нових даних та тим самим робити модель більш адаптивною.

Інтерпретація отриманих результатів забезпечується за рахунок застосування методів Explainable AI: використання LIME, SHAP або інших подібних технологій допомагає виявити, які ознаки сприяли рішенням моделі, і на яких паттернах вона зосереджується. Це важливо, оскільки користувачі модельних результатів (менеджери, інвестори, консультанти) часто потребують не лише

чіткої цифри ризику, а й розуміння того, які фактори впливають на цей ризик. Таким чином, прозорість є ключовим елементом ефективного використання розробленого методу в реальних умовах.

Практичне використання цього комплексного підходу може істотно покращити управління ризиками у стартап-середовищі. Аналітики отримують інструмент, що дає змогу швидко оцінити перспективи різних проєктів, побачити потенційні "слабкі місця" та розуміти, як певні заходи (наприклад, збільшення інвестицій, найм технічних спеціалістів, зміна стратегій виходу на ринок) можуть вплинути на рівень ризиків. Динамічний характер моделі дозволяє оновлювати прогнози при зміні ринкових умов, появі нових трендів або отриманні нових даних про діяльність стартапу.

3.2 Опис роботи алгоритма

У світі сучасних технологій, де дані є важливим ресурсом, ефективне управління ризиками в стартапах та інших бізнес-проєктах набуває все більшої важливості. Одним із найбільш дієвих способів оцінки потенційних загроз є використання алгоритмів, які дозволяють автоматизувати процес оцінки рівня ризику. Цей процес включає в себе використання машинного навчання для класифікації стартапів за рівнями ризику на основі великої кількості змінних, таких як фінансові показники, характеристика команди, ринкові умови та інші фактори, що можуть впливати на успішність бізнесу. Однак для того, щоб алгоритм міг правильно класифікувати стартапи за рівнем ризику, дані повинні бути попередньо оброблені, нормалізовані та зменшені в розмірності, що дозволяє зменшити складність і покращити ефективність моделі.

Використання такої методології дає можливість підвищити точність прогнозів і допомогти інвесторам або підприємцям у виборі стартапів з оптимальним рівнем ризику для інвестування або розвитку. Застосування алгоритмів класифікації та прогнозування дозволяє не лише виявити потенційно високоризикові стартапи, але й запропонувати конкретні рекомендації для

зниження цього ризику на підставі виявлених закономірностей. Описаний алгоритм використовує популярні методи, такі як KMeans для класифікації та метод головних компонент (PCA) для зменшення розмірності даних, що дозволяє підвищити точність результатів без зайвих обчислювальних витрат.

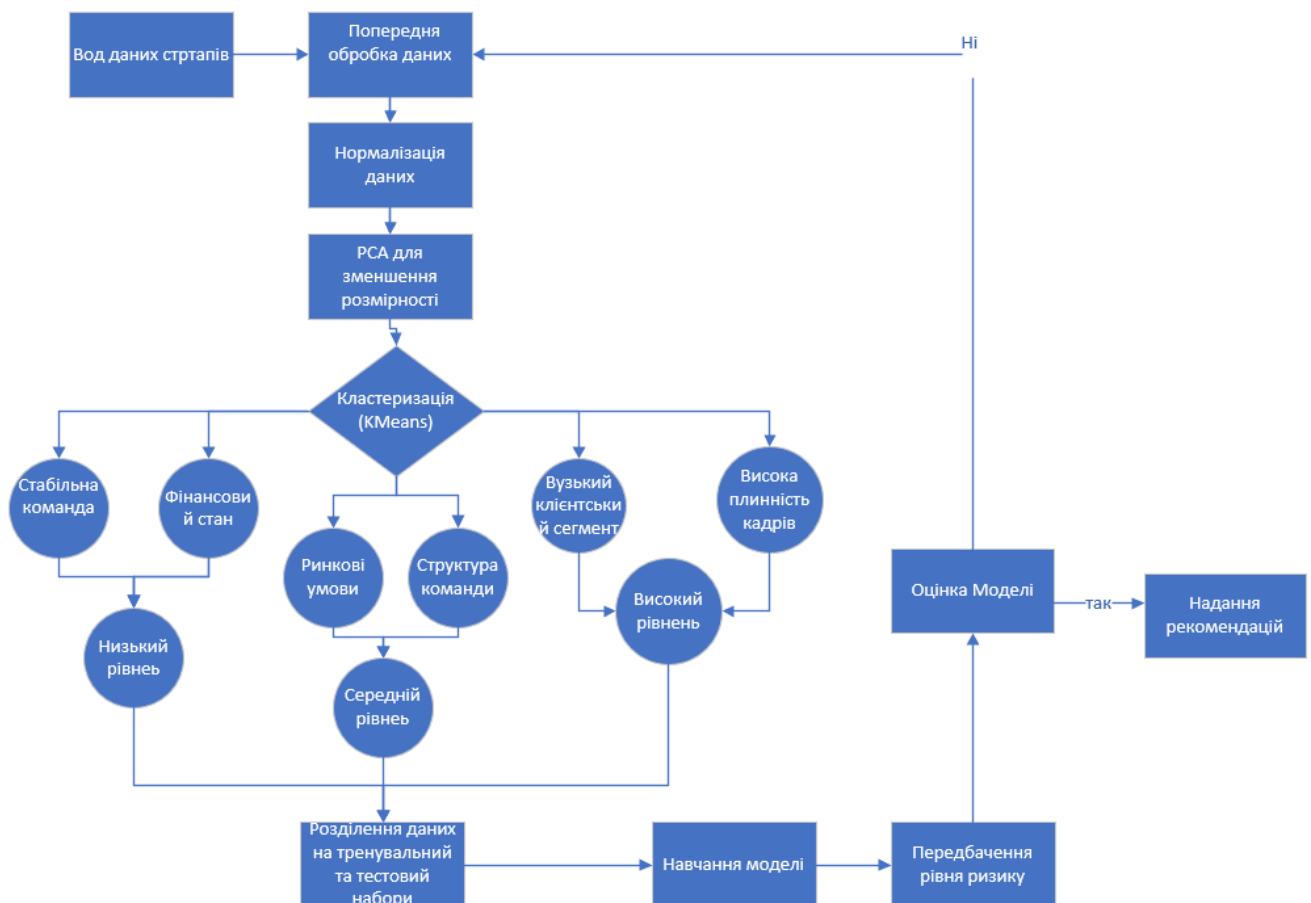


Рис. 1.1 Алгоритм прогнозування ризиків

Першим етапом є підготовка даних, яка включає їх очищення від помилок, аномалій та неповних значень. Це важливий крок, оскільки помилки в даних можуть серйозно вплинути на точність моделей машинного навчання. Після цього проводиться нормалізація даних, щоб усі змінні мали однаковий масштаб. Це забезпечує рівноправний вплив усіх змінних на результати класифікації, не даючи переваги жодному з факторів тільки через більші числові значення. Нормалізація даних є ключовою для створення збалансованої моделі, яка забезпечить коректні

результати для всіх типів вхідних даних.

На наступному етапі застосовується метод головних компонент (PCA), який дозволяє зменшити розмірність даних. У разі, коли набір даних містить велику кількість змінних, це може призвести до складності у моделюванні і, як результат, до зниження ефективності алгоритму. Зменшення розмірності за допомогою PCA дозволяє зберегти важливу інформацію, зменшивши кількість змінних, що використовуються при класифікації. Це не лише прискорює процес навчання моделі, але й підвищує її точність, оскільки модель фокусується на найбільш значущих характеристиках стартапу, зменшуючи вплив неважливих або корельованих факторів.

Після зменшення розмірності даних застосовується класифікація за допомогою алгоритму KMeans, який класифікує стартапи на основі їхніх характеристик за рівнями ризику. Класифікація здійснюється з урахуванням фінансових показників, ринкових умов, структури команди тощо. Стартапи можуть бути розподілені на три категорії: низький, середній та високий рівень ризику. Цей етап дозволяє виявити стартапи, що потребують додаткової уваги, а також допомагає інвесторам зосередити свої ресурси на проєктах з найвищим потенціалом для росту.

Проте для підвищення точності класифікації та прогнозування, особливо коли мова йде про складніші задачі, на допомогу приходять потужніші методи машинного навчання, такі як RandomForestClassifier. Це ансамблевий метод, який базується на комбінуванні декількох дерев рішень для досягнення високої точності та зниження ймовірності помилок. RandomForestClassifier особливо ефективний у випадках, коли існує велика кількість змінних, оскільки він не схильний до перенавчання і здатний обробляти великий обсяг даних без значної втрати якості.

Навчання моделі з використанням RandomForestClassifier є наступним етапом. На цьому етапі модель RandomForestClassifier навчається на тренувальних даних, які містять характеристики стартапів і відповідні їм рівні ризику. Алгоритм RandomForest створює кілька дерев рішень, кожне з яких приймає рішення щодо

класифікації стартапу, враховуючи різні фактори.

Важливою перевагою цього методу є те, що кожне дерево в лісі робить класифікацію незалежно від інших дерев, і в кінцевому підсумку всі дерева голосують за найкращий варіант. Така стратегія дозволяє значно зменшити ймовірність помилок, оскільки помилки одного дерева можуть бути компенсовані іншими деревами в лісі.

У процесі навчання моделі важливо налаштувати кілька параметрів `RandomForestClassifier`, таких як кількість дерев у лісі (`n_estimators`) і максимальна глибина дерева (`max_depth`), щоб досягти оптимального балансу між точністю і швидкістю моделювання. Оскільки `RandomForestClassifier` є ансамблевим методом, він менш схильний до перенавчання порівняно з окремими деревами рішень. Це робить його ідеальним для класифікації стартапів, оскільки він може адаптуватися до складних даних з багатьма змінними та зберігати високу точність прогнозів навіть у випадку великої кількості вихідних параметрів.

Далі модель проходить через етап перевірки на тестових даних. Після того як `RandomForestClassifier` навчається на тренувальних даних, його точність оцінюється за допомогою тестових даних, які не використовувалися під час навчання.

Це дозволяє перевірити, як добре модель може узагальнювати інформацію та класифікувати нові стартапи. Для цього використовуються метрики, такі як точність, повнота та F1-міра, які допомагають об'єктивно оцінити якість роботи моделі.

Після того як модель пройшла навчання та тестування, вона готова до прогнозування рівня ризику для нових стартапів. Модель `RandomForestClassifier` класифікує стартапи за рівнем ризику на основі отриманих від тренувальних даних закономірностей, що дає змогу точно визначити ризики для кожного стартапу.

Завершальним етапом є надання рекомендацій для зменшення рівня ризику для стартапів з високим рівнем ризику. Модель може порекомендувати коригування у структурі команди, фінансових стратегіях чи маркетингових

підходах для мінімізації можливих загроз. Застосування RandomForestClassifier у поєднанні з іншими методами, такими як PCA та KMeans, дозволяє створити потужний інструмент для управління ризиками в стартапах, що дає змогу знизити ймовірність фінансових втрат і підвищити ефективність інвестицій.

3.3 Опис використаних програмних засобів

У сучасному світі, де обсяги даних зростають експоненціально, ефективний аналіз та обробка інформації стають критично важливими для прийняття обґрунтованих рішень у різних галузях науки, бізнесу та промисловості. Для досягнення високих результатів у рамках дипломної роботи було обрано кілька потужних програмних засобів, серед яких основну роль відіграють мова програмування Python та її ключові бібліотеки: NumPy, Pandas та scikit-learn.

Ці інструменти забезпечують необхідну гнучкість, продуктивність та точність для виконання складних обчислювальних завдань, аналізу великих обсягів даних та побудови надійних моделей машинного навчання. У цьому розділі детально розглянемо кожен із цих програмних засобів, їхні особливості, переваги та роль у виконанні дипломної роботи.

Python є високорівневою, інтерпретованою мовою програмування, яка набуває все більшої популярності завдяки своїй універсальності, простоті та широким можливостям. Python використовується у різних сферах, таких як наука про дані, машинне навчання, веб-розробка, автоматизація та багато інших. Основні переваги Python включають:

Простота та читабельність синтаксису: Python розроблений з акцентом на читабельність коду, що дозволяє розробникам швидко писати та розуміти програмний код. Це особливо важливо у наукових дослідженнях, де швидкість розробки та ефективність роботи мають вирішальне значення.

Відкрите джерело та безкоштовність: Python є мовою з відкритим вихідним кодом, що забезпечує безкоштовний доступ до його функціоналу. Це робить його доступним для широкого кола користувачів та сприяє активному розвитку

спільноти.

Велика кількість бібліотек та фреймворків: Python має величезну екосистему бібліотек, що покривають широкий спектр завдань від обробки даних до розробки веб-додатків. Це дозволяє використовувати готові рішення та інструменти для швидкої реалізації проектів.

Активна спільнота користувачів: Велика та активна спільнота розробників забезпечує доступ до численних ресурсів, таких як документація, навчальні матеріали, форуми та приклади коду. Це спрощує процес навчання та вирішення виникаючих проблем.

Інтерпретована мова: Python є інтерпретованою мовою, що дозволяє швидко тестувати та запускати код без необхідності компіляції. Це сприяє більш швидкому циклу розробки та налагодження програм.

Міжплатформенність: Python працює на різних операційних системах, таких як Windows, macOS та Linux, що робить його універсальним інструментом для розробки та аналізу даних.

Гнучкість та універсальність: Python використовується у різних галузях, таких як веб-розробка, автоматизація, наука про дані, фінанси, що робить його універсальним інструментом для різноманітних завдань.

Python широко використовується для аналізу даних, автоматизації рутинних завдань, розробки веб-додатків, створення скриптів для обробки тексту та числових даних, розробки ігор, а також у наукових дослідженнях для моделювання та симуляцій. Наприклад, у фінансах Python може використовуватись для аналізу ринкових тенденцій, у медицині — для обробки медичних зображень, а в маркетингу — для аналізу поведінки споживачів.

NumPy (Numerical Python) є фундаментальною бібліотекою для наукових обчислень у Python. Вона забезпечує підтримку великих багатовимірних масивів та матриць, а також широкий спектр математичних функцій для їх обробки. NumPy є базовим інструментом для багатьох інших бібліотек Python, таких як Pandas та scikit-learn, оскільки надає ефективні структури даних та функції для роботи з числовими даними.

NumPy широко використовується у різних сферах, таких як фізика, інженерія, фінанси та біоінформатика. Наприклад, у фізиці NumPy може застосовуватись для моделювання фізичних процесів, таких як рух тіл або розподіл температури.

У фінансах NumPy допомагає аналізувати фінансові дані та розраховувати ризики інвестиційних портфелів. У біоінформатиці NumPy використовується для обробки геномних даних та аналізу біологічних процесів. Завдяки своїй ефективності та гнучкості, NumPy стає незамінним інструментом для виконання числових обчислень та обробки даних у різних наукових та технічних завданнях.

Pandas є однією з найпопулярніших бібліотек для обробки та аналізу даних у Python. Вона надає зручні структури даних, такі як DataFrame та Series, що дозволяють ефективно працювати з табличними даними. Pandas є потужним інструментом для виконання різноманітних операцій з даними, включаючи фільтрацію, агрегацію, трансформацію та візуалізацію інформації.

Pandas широко використовується у різних сферах, таких як економіка, соціологія, екологія та медицина. Наприклад, у економіці Pandas може використовуватись для аналізу макроекономічних показників, таких як ВВП, інфляція та безробіття.

Таблиця 3.1

Порівняльні характеристики бібліотек NumPy, Pandas та scikit-learn

Бібліотека	Основні функції	Переваги	Обмеження
NumPy	Робота з багатовимірними масивами, математичні функції	Висока продуктивність, фундаментальна	Мала кількість вбудованих функцій для аналізу

Продовження таблиці 3.1

Порівняльні характеристики бібліотек NumPy, Pandas та scikit-learn

Бібліотека	Основні функції	Переваги	Обмеження
------------	-----------------	----------	-----------

Pandas	Обробка та аналіз табличних даних, імпорт/експорт даних	Зручні структури даних (DataFrame), багатий набір функцій	Вимагає знань з обробки даних
scikit-learn	Машинне навчання, класифікація, регресія, кластеризація	Великий набір алгоритмів, зручний інтерфейс	Обмежені можливості для глибокого навчання

У соціології Pandas допомагає аналізувати опитувальні дані, вивчати тенденції та взаємозв'язки між різними соціальними факторами. У медицині Pandas застосовується для обробки та аналізу даних пацієнтів, вивчення ефективності лікарських препаратів та дослідження епідеміологічних тенденцій. Завдяки своїй гнучкості та потужності, Pandas є незамінним інструментом для виконання комплексного аналізу даних у різних наукових та професійних дослідженнях.

scikit-learn є однією з найпотужніших та найпопулярніших бібліотек для машинного навчання та статистичного аналізу у Python. Вона надає широкий набір інструментів для виконання різноманітних завдань, включаючи класифікацію, регресію, кластеризацію, зменшення розмірності та оцінку моделей. scikit-learn є надзвичайно важливим інструментом для побудови та впровадження моделей машинного навчання завдяки своїй простоті використання та високій ефективності.

Основні характеристики та переваги scikit-learn:

Різноманітність алгоритмів: scikit-learn підтримує велику кількість алгоритмів машинного навчання, включаючи лінійну та логістичну регресію, дерева рішень, випадкові ліси, підтримувальні вектори машин (SVM), кластеризацію K-середніх, методи зменшення розмірності (PCA, t-SNE) та багато інших. Це дозволяє вирішувати різні типи завдань без необхідності впровадження алгоритмів з нуля.

Зручний інтерфейс: Бібліотека має єдиний та зрозумілий інтерфейс для

різних методів, що спрощує її використання. Завдяки цьому користувачі можуть легко переключатися між різними алгоритмами та налаштовувати їхні параметри.

Інтеграція з іншими бібліотеками: `scikit-learn` легко інтегрується з `Pandas` та `NumPy`, що забезпечує безшовну роботу з даними. Це дозволяє ефективно комбінувати різні етапи аналізу даних, від їхньої підготовки до побудови та оцінки моделей машинного навчання.

Висока ефективність та масштабованість: `scikit-learn` оптимізований для роботи з великими наборами даних, що робить його ідеальним інструментом для обробки великих обсягів інформації без значних втрат продуктивності.

Документація та підтримка: Бібліотека має детальну документацію та численні приклади використання, що полегшує процес навчання та вирішення виникаючих проблем. Крім того, `scikit-learn` має активну спільноту користувачів, що забезпечує доступ до додаткових ресурсів та підтримки.

Методи оцінки моделей: `scikit-learn` надає різноманітні метрики для оцінки якості моделей, такі як точність, повнота, F-мера для класифікації, середнє квадратичне відхилення (MSE) для регресії та інші, що дозволяє об'єктивно оцінювати ефективність моделей.

Вибір відповідних програмних засобів є ключовим фактором успіху будь-якого дослідницького проекту, зокрема дипломної роботи. Мова програмування Python разом із бібліотеками `NumPy`, `Pandas` та `scikit-learn` надають потужний та гнучкий інструментарій для збору, обробки, аналізу та моделювання даних. Ці інструменти дозволяють ефективно працювати з великими обсягами інформації, виконувати складні обчислювальні завдання та будувати надійні моделі машинного навчання, що є основою сучасного аналізу даних та наукових досліджень.

ВИСНОВКИ

У ході виконання магістерської роботи проведено комплексний аналіз існуючих підходів до оцінки та прогнозування ризиків у стартап-проектах, що дозволило не лише глибше зрозуміти сучасний стан проблеми, але й розробити нову методологію для вдосконалення управління ризиками в умовах високої невизначеності. У сучасному світі стартапи стикаються з численними викликами, такими як нестача стабільних часових рядів, невизначеність ринкової кон'юнктури, швидкі зміни зовнішнього середовища та обмеженість ресурсів. У цих умовах традиційні методи аналізу ризиків, такі як регресійний аналіз, аналіз дерев рішень, сценарний аналіз чи імітаційне моделювання, виявляються недостатньо ефективними, оскільки вони не враховують багатовимірність, динамічність даних і специфіку інноваційних проєктів. Це зумовило необхідність розробки нових підходів, які інтегрують сучасні технології машинного навчання та інтелектуального аналізу даних.

У межах виконаної роботи запропоновано багаторівневу методологію аналізу ризиків, яка поєднує переваги байєсових моделей, класичних методів статистики, алгоритмів машинного навчання. На першому рівні застосовано байєсові моделі, які дозволяють інтегрувати експертні знання з наявними даними. Це забезпечує можливість динамічного оновлення прогнозів залежно від змін у системі або появи нових даних. Байєсовий підхід є особливо цінним у ситуаціях, коли даних недостатньо, а також у контексті нових стартапів, де історичних даних майже немає. Завдяки цьому підходу вдалося створити первинну модель для оцінки ризиків, яка слугує базою для подальших досліджень.

Другий рівень запропонованої методології базується на алгоритмах машинного навчання, таких як логістична регресія, алгоритми опорних векторів (SVM), ансамблеві методи (Random Forest, Gradient Boosting). Ці алгоритми забезпечують автоматизацію процесу аналізу великих обсягів

даних, виявлення прихованих закономірностей та побудову моделей з високою точністю. Зокрема, застосування ансамблевих методів дозволило мінімізувати ризик перенавчання та підвищити надійність прогнозів.

Логістична регресія забезпечила інтерпретованість моделей на початкових етапах, тоді як SVM із нелінійними ядрами надала можливість аналізувати складні багатовимірні залежності, що недоступні для традиційних методів.

На третьому рівні використано глибокі нейронні мережі, включаючи згорткові мережі (CNN), рекурентні мережі (RNN, LSTM, GRU) та трансформери. Ці моделі продемонстрували високу ефективність при роботі зі складними, неструктурованими та різнорідними даними, такими як текстові документи, часові ряди, зображення чи сенсорні дані. Завдяки здатності нейронних мереж до автоматичного вилучення ознак із сирих даних вдалося досягти значного підвищення точності прогнозів ризиків. Зокрема, трансформери забезпечили ефективну обробку текстової інформації, дозволяючи інтегрувати такі джерела даних, як звіти, новини, фінансові документи, у загальну систему аналізу ризиків.

У рамках роботи було створено програмне забезпечення, яке реалізує запропоновану методологію. Для його розробки використано мову програмування Python та бібліотеки машинного навчання, такі як scikit-learn, TensorFlow, Keras. Програмне забезпечення забезпечує інтеграцію даних із різних джерел, їхню попередню обробку, аналіз та прогнозування ризиків. Зокрема, програмний продукт дозволяє автоматично оцінювати рівень ризику стартапу, виявляти критичні фактори, що впливають на ризик, та надавати рекомендації щодо його зниження.

Результати магістерської роботи підтвердили перспективність застосування машинного навчання та сучасних алгоритмів аналізу даних для прогнозування ризиків у стартапах. Розроблена методологія може знайти широке застосування не лише у сфері стартапів, але й у фінансах, охороні здоров'я, промисловості та інших галузях. Запропоноване рішення є

універсальним, масштабованим та адаптивним, що робить його цінним інструментом для прийняття стратегічних рішень і управління ризиками в умовах невизначеності.

Результати дослідження апробовано шляхом публікації тез доповідей на конференціях:

1. Дзядевич Д.Д., Садовенко В.С. Аналіз і прогнозування ризиків у стартапах за допомогою сучасних ML-алгоритмів // Всеукраїнська науково-технічна конференція "Виклики та рішення в програмній інженерії".- Київ: ДУІКТ, 2014 – С. 100-102

2. Дзядевич Д.Д., Садовенко В.С. Автоматизація оцінки ризиків стартапів за допомогою методів штучного інтелекту // Всеукраїнська науково-технічна конференція "Виклики та рішення в програмній інженерії".- Київ: ДУІКТ, 2024 – С. 157-159.

ПЕРЕЛІК ПОСИЛАНЬ

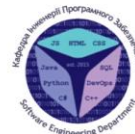
1. Ries, Eric[1]. *The Lean Startup: How Today's Entrepreneurs Use Continuous Innovation to Create Radically Successful Businesses*. Crown Business, 2011.
2. Kahneman, Daniel[2]. *Thinking, Fast and Slow*. Farrar, Straus and Giroux, 2011.
3. Grant, Adam[3]. *Originals: How Non-Conformists Move the World*. Viking, 2016.
4. Christensen, Clayton M[4]. *The Innovator's Dilemma: When New Technologies Cause Great Firms to Fail*. Harvard Business Review Press, 1997.
5. Christensen, Clayton M. *The Innovator's Solution: Creating and Sustaining Successful Growth*. Harvard Business Review Press, 2003.
6. Christensen, Clayton M. *Competing Against Luck: The Story of Innovation and Customer Choice*. Harper Business, 2016.
7. Feld, Brad. *Venture Deals: Be Smarter Than Your Lawyer and Venture Capitalist*. Wiley, 2012.
8. Feld, Brad. *Startup Communities: Building an Entrepreneurial Ecosystem in Your City*. Wiley, 2012.
9. Feld, Brad. *The Startup Community Way: Evolving an Entrepreneurial Ecosystem*. Wiley, 2020.
10. Семешкін А. М. Модель оцінки ризиків у проєктному менеджменті з використанням машинного навчання. // *Економічний вісник*. – 2023. – №4. – С. 45-58.
11. Михайлов Н. О. Адаптивна модель планування проєктів та оцінки ризиків із використанням машинного навчання. // *Вісник Ужгородського національного університету*. – 2024. – №2. – С. 123-135.
12. Грас Дж. Data Science. Наука про дані з нуля. // *Видавництво "Наука і Техніка"*. – 2022. – 320 с.
13. Кононова К. Ю. Машинне навчання: методи та моделі : підручник для студентів вищих навчальних закладів. // *Одеса: ОНУ ім. І. І. Мечникова*. – 2021. – 450 с.
14. Жадан Т. В. Аналіз та прогнозування ризиків в діяльності підприємства. // *Навчально-методичний посібник*. – 2021. – 150 с.
15. Керзнер Г. Управління проєктами: Системний підхід до планування, розкладання та контролю. // *Видавництво "Project Management Institute"*. – 2017. – 850 с.
16. Штуб А., Бард Дж. Ф., Глоберсон С. Управління проєктами: Процеси, методології та економіка. // *Видавництво "Pearson Prentice Hall"*. – 2005. – 600 с.
17. Хіллсон Д. Управління ризиками в проєктах. // *Видавництво "Routledge"*. – 2017. – 250 с.
18. Раз Т., Шенхар А. Аналіз ризиків у проєктному менеджменті: Методи та інструменти. // *Видавництво "Wiley"*. – 2002. – 300 с.
19. Гудфеллоу Я., Бенджіо Й., Курвіль А. Глибоке навчання. // *Видавництво "MIT Press"*. – 2016. – 800 с.

20. Бішоп К. Розпізнавання візерунків та машинне навчання. // *Видавництво "Springer"*. – 2006. – 738 с.
21. Нг Е. Прагнення машинного навчання. // *Видавництво "O'Reilly Media"*. – 2019. – 400 с.
22. Мюллер А. К., Гуїдо С. Вступ до машинного навчання з Python. // *Видавництво "O'Reilly Media"*. – 2016. – 400 с.
23. Шолле Ф. Глибоке навчання з Python. // *Видавництво "Manning Publications"*. – 2017. – 384 с.
24. Янсенс Дж. Data Science у командному рядку. // *Видавництво "O'Reilly Media"*. – 2014. – 212 с.
25. Келлехер Дж. Д., Мак Неймі Б., Д'Арсі А. Введення в машинне навчання. // *Видавництво "MIT Press"*. – 2015. – 624 с.
26. Шалев-Шварц Ш., Бен-Давід Ш. Розуміння машинного навчання: Від теорії до алгоритмів. // *Видавництво "Cambridge University Press"*. – 2014. – 410 с.
27. Картера Д. Штучний інтелект. Машинне навчання. // *Видавництво "Book-Shop"*. – 2023. – 350 с.
28. Тардіф А. 6 найкращих книг усіх часів про машинне навчання та ШІ. // *Unite.AI*. – 2024.
29. Redmon, J., & Farhadi, A. (2018). *YOLOv3: An Incremental Improvement*. arXiv preprint arXiv:1804.02767.
30. Ren, S., He, K., Girshick, R., & Sun, J. (2017). *Faster R-CNN: Towards Real-Time Object Detection with Region Proposal Networks*. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 39(6), 1137-1149.
31. Lin, T. Y., Goyal, P., Girshick, R., He, K., & Dollár, P. (2017). *Focal Loss for Dense Object Detection*. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 42(2), 318-327.

ДОДАТОК А. ДЕМОНСТРАЦІЙНІ МАТЕРІАЛИ



ДЕРЖАВНИЙ УНІВЕРСИТЕТ
ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ
НАВЧАЛЬНО-НАУКОВИЙ ІНСТИТУТ ІНФОРМАЦІЙНИХ
ТЕХНОЛОГІЙ



КАФЕДРА ІНЖЕНЕРІЇ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ

Магістерська робота

«Методика аналізу та прогнозування ризиків у стартап-проектах на основі технології машинного навчання»

Виконав: студент групи ПДМ-62 Дзядевич Дмитро Дмитрович

Керівник: к.ф.-м.н.професор кафедри ІПЗ Садовенко Володимир Сергійович

Київ - 2025

МЕТА, ОБ'ЄКТА ТА ПРЕДМЕТ ДОСЛІДЖЕННЯ

Мета роботи: Оптимізація процесу аналізу і прогнозування ризиків у стартап-проектах за рахунок використання методів машинного навчання, що забезпечують ефективну оцінку рівня ризику та надання рекомендацій для їх мінімізації.

Об'єкт дослідження: Процес аналізу і прогнозування ризиків у стартап-проектах.

Предмет дослідження: Методи машинного навчання та алгоритми кластеризації, які використовуються для аналізу і прогнозування ризиків у стартапах.

Актуальність роботи

Метод	Переваги	Недоліки
Методи статистичного аналізу	Простота реалізації.	Обмежена адаптивність до змін
Регресійний аналіз	Здатність кількісно оцінювати вплив факторів.	Обмежена застосовність для нелінійних залежностей
Кластеризація (K-Means)	Ефективна для групування об'єктів	Складнощі при роботі з неоднорідними даними
Алгоритми опорних векторів (SVM)	Стійкість до шуму у даних.	Складність налаштування ядер
Ансамблеві методи (Random Forest, Gradient Boosting)	Висока точність та стійкість до перенавчання	Необхідність налаштування великої кількості параметрів

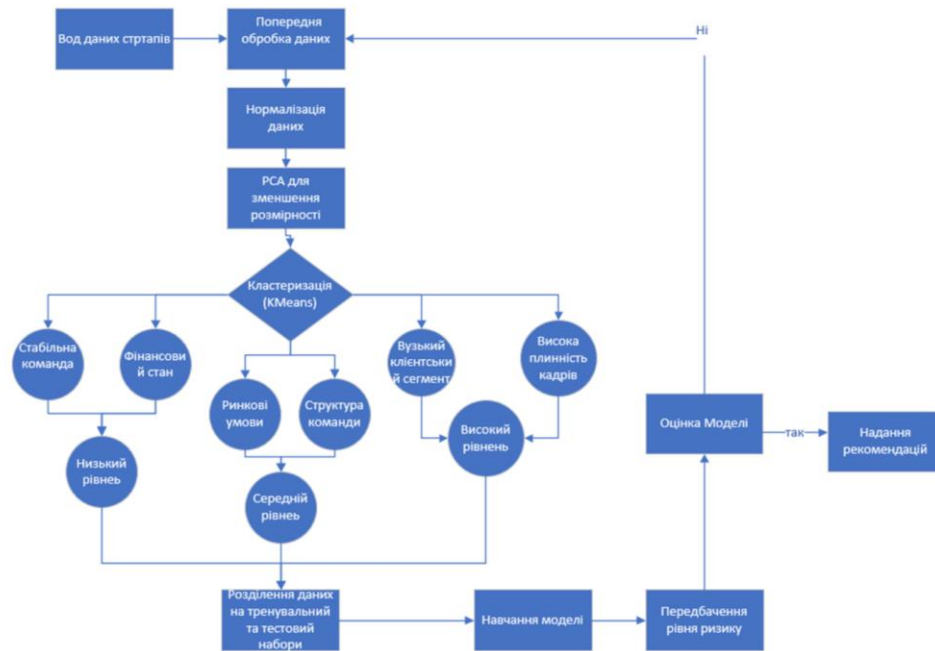
3

Об'єкт кластеризації

Категорія об'єктів	Характеристики об'єктів	Мета кластеризації
Фінансові показники	Рівень доходів, витрати, інвестиції, фінансова стабільність, прогнозований грошовий потік	Виділення стартапів із високим фінансовим ризиком або потенціалом фінансового успіху
Команда стартапу	Кількість учасників, професійний досвід, рівень кваліфікації, ступінь залученості	Оцінка впливу якості команди на рівень ризику
Ринкові умови	Розмір ринку, конкурентоспроможність, доступ до клієнтів, рівень попиту	Аналіз ринкових факторів, що впливають на ризики стартапу
Ризики зовнішнього середовища	Економічна стабільність, регуляторні вимоги, політичні ризики, макроекономічні фактори	Оцінка зовнішніх загроз, що впливають на можливість успіху стартапу

5

Алгоритм прогнозування ризиків



Програмні засоби



ПРАКТИЧНИЙ РЕЗУЛЬТАТ

```
Звіт про класифікацію:
      precision    recall  f1-score   support

     0       0.50      1.00      0.67         1
     1       1.00      0.67      0.80         3

 accuracy               0.75         4
 macro avg              0.75      0.83      0.73         4
weighted avg              0.88      0.75      0.77         4

Стартап 1: Рівень ризику - Середній, Рекомендації - Покращити маркетингову стратегію
Стартап 2: Рівень ризику - Низький, Рекомендації - Підтримувати поточний рівень
Стартап 3: Рівень ризику - Високий, Рекомендації - Збільшити фінансування та зменшити кон
Стартап 4: Рівень ризику - Середній, Рекомендації - Покращити маркетингову стратегію
Стартап 5: Рівень ризику - Низький, Рекомендації - Підтримувати поточний рівень
Стартап 6: Рівень ризику - Високий, Рекомендації - Збільшити фінансування та зменшити кон
Стартап 7: Рівень ризику - Низький, Рекомендації - Підтримувати поточний рівень
Стартап 8: Рівень ризику - Середній, Рекомендації - Покращити маркетингову стратегію
Стартап 9: Рівень ризику - Високий, Рекомендації - Збільшити фінансування та зменшити кон
Стартап 10: Рівень ризику - Низький, Рекомендації - Підтримувати поточний рівень
```

8

Порівняльний аналіз

Порівняльний Аналіз

	Критерій	Існуючі методи	Запропонований метод
1	Точність прогнозування	Аналіз дерев рішень, Регресійний аналіз	85%
2	Адаптивність	Сценарний аналіз	5 балів
3	Вимоги до ресурсів	Аналіз часових рядів	Середні
4	Універсальність	Імітаційне моделювання	Висока

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

8

ВИСНОВКИ

- Проаналізовано сучасні методи прогнозування ризиків у стартап-проектах, що базуються на використанні машинного навчання, статистичних підходів та алгоритмів глибокого навчання.
- Досліджено приклади застосування технологій машинного навчання у прогнозуванні ризиків
- Вдосконалено підхід до прогнозування ризиків завдяки інтеграції машинного навчання, байєсових методів та ансамблевих алгоритмів

10

ПУБЛІКАЦІЇ ТА АПРОБАЦІЯ РОБОТИ

Тези:

1. Дзядевич Д.Д., Садовенко В.С. Аналіз і прогнозування ризиків у стартапах за допомогою сучасних ML-алгоритмів // Всеукраїнська науково-технічна конференція "Виклики та рішення в програмній інженерії".- Київ: ДУІКТ, 2014 – С. 100-102

2. Дзядевич Д.Д., Садовенко В.С. Автоматизація оцінки ризиків стартапів за допомогою методів штучного інтелекту // Всеукраїнська науково-технічна конференція "Виклики та рішення в програмній інженерії".- Київ: ДУІКТ, 2024 – С. 157-159.

11

ДЯКУЮ ЗА УВАГУ!

ДОДАТОК Б. ЛІСТИНГИ ПРОГРАМНИХ МОДУЛІВ

```
import numpy as np
import pandas as pd
from sklearn.preprocessing import
StandardScaler
from sklearn.decomposition import PCA
from sklearn.cluster import KMeans
from sklearn.model_selection import
train_test_split
from sklearn.ensemble import
RandomForestClassifier
from sklearn.metrics import
classification_report, accuracy_score

df = pd.read_csv('startups.csv')
df = df.dropna()

features_for_clustering = [
    'team_stability',
    'financial_status',
    'market_conditions',
    'client_segment',
    'staff_turnover'
]

X = df[features_for_clustering].values
scaler = StandardScaler()
X_scaled = scaler.fit_transform(X)
pca = PCA(n_components=2)
X_pca = pca.fit_transform(X_scaled)
kmeans = KMeans(n_clusters=3,
random_state=42)
kmeans_labels =
kmeans.fit_predict(X_pca)
df['cluster_risk'] = kmeans_labels
cluster_to_risk = {0: 'Низький рівень',
1: 'Середній рівень', 2: 'Високий
рівень'}
df['cluster_risk'] =
df['cluster_risk'].map(cluster_to_risk)
df['risk_num'] =
df['cluster_risk'].map({'Низький
рівень': 0, 'Середній рівень': 1,
'Високий рівень': 2})
y = df['risk_num'].values

X_final = X_scaled
X_train, X_test, y_train, y_test =
train_test_split(X_final, y, test_size=0.2,
random_state=42)
model =
RandomForestClassifier(n_estimators=1
00, max_depth=None,
random_state=42)
model.fit(X_train, y_train)
y_pred = model.predict(X_test)
print("Accuracy Score:",
accuracy_score(y_test, y_pred))
print("\nClassification Report:\n",
classification_report(y_test, y_pred))
test_predictions =
pd.DataFrame({'PredictedRisk':
y_pred})
high_risk_indices =
test_predictions[test_predictions['Predict
edRisk'] == 2].index
for idx in high_risk_indices:
    print(f"Стартап {idx}: Рекомендації
для зменшення ризику.")
```