

**НАЦІОНАЛЬНИЙ ТЕХНІЧНИЙ УНІВЕРСИТЕТ УКРАЇНИ
«КИЇВСЬКИЙ ПОЛІТЕХНІЧНИЙ ІНСТИТУТ
імені ІГОРЯ СІКОРСЬКОГО»**

факультет інформатики та обчислювальної техніки
(повна назва інституту/факультету)

кафедра автоматика та управління в технічних системах
(повна назва кафедри)

«На правах рукопису»
УДК _____

«До захисту допущено»

Завідувач кафедри
_____ О. І. РОЛІК
(підпис) (ініціали, прізвище)

“ ” _____ 2018
р.

Магістерська дисертація

зі спеціальності (спеціалізації) 126 «Інформаційні системи та технології»
(код і назва спеціальності)

на тему: Система прогнозування фінансових ринків із використанням рекурентних нейромереж

Виконав : студент 6 курсу, групи ІА-72мп
(шифр групи)

Прохоров Дмитро Павлович _____
(прізвище, ім'я, по батькові) (підпис)

Науковий керівник доцент, к.тех.н. Кравець П.І. _____
(посада, науковий ступінь, вчене звання, прізвище та ініціали) (підпис)

Консультант _____
(назва розділу) (науковий ступінь, вчене звання, прізвище, ініціали) (підпис)

Рецензент _____
(посада, науковий ступінь, вчене звання, науковий ступінь, прізвище та ініціали) (підпис)

Засвідчую, що у цій магістерській дисертації немає запозичень з праць інших авторів без відповідних посилань.

Студент _____
(підпис)

Київ – 2018 рік

**Національний технічний університет України
“Київський політехнічний інститут
імені Ігоря Сікорського”**

Факультет інформатики та обчислювальної техніки
(повна назва)

Кафедра автоматики та управління в технічних системах
(повна назва)

Ступінь вищої освіти –другий (магістерський)
(код, назва)

Спеціальність 126 «Інформаційні системи та технології»
(код, назва)

ЗАТВЕРДЖУЮ

Завідувач кафедри

(підпис) О. І. РОЛІК
(ініціали, прізвище)

“ ____ ” _____ 2018_р.

ЗАВДАННЯ

на магістерську дисертацію студенту

Прохорову Дмитру Павловичу

(прізвище, ім'я, по батькові)

1.Тема дисертації Система прогнозування фінансових ринків із використанням рекурентних нейромереж

Науковий керівник дисертації Кравець Петро Іванович, к. т. н.,
доцент

затверджені наказом по університету від “ 29 ” жовтня 2018_р.№ _____

2.Строк подання студентом дисертації “ 4 ” грудня 2018_р.

3. Об`єкт дослідження: нейронні мережі із використанням рекурентних зв'язків

4.Зміст пояснювальної записки: а) призначення та область застосування;
б) розробка структурних схем системи та нейромереж; в) вибір окремих вузлів;
г) розробка діаграми класів; д) розробка алгоритмів обробки фінансових новин; е) моделювання за допомогою пакетів мови Python

5. Перелік графічного (ілюстративного) матеріалу

Структурна схема системи, функціональна схема системи, діаграма варіантів використання, структурна схема модулю обробки фінансових новин, алгоритм роботи модулю фінансових новин, діаграма прецедентів системи, діаграма розгортання системи.

6. Консультанти розділів проекту:

Розділ	Прізвище, ініціали та посада консультанта	Підпис, дата	
		Завдання видав	Завдання прийняв

7. Дата видачі завдання “_29_” _жовтня_ 2018_р.

Календарний план

№ з/п	Назва етапів виконання магістерської дисертації	Термін виконання етапів проекту	Примітка
1	Огляд існуючих рішень та розробка структури системи	10.11.2018	
2	Вибір окремих вузлів	15.11.2018	
3	Розробка структурної схеми системи	20.11.2018	
4	Розробка схеми алгоритму роботи та модулю обробки фінансових новин	25.11.2018	
5	Моделювання за допомогою засобів Python	29.11.2018	
6	Оформлення текстової та графічної документації	2.12.2018	
7	Представлення до захисту	4.12.2018	

Студент

(підпис)

Прохоров Д.П.

(ініціали, прізвище)

Керівник проекту

(підпис)

Кравець П.І.

(ініціали, прізвище)

РЕФЕРАТ

Магістерська дисертація: 125 с., 21 рис., 25 табл., 8 додатків, 34 джерела.

Тема магістерської дисертації “Прогнозування фінансових ринків із використанням рекурентних нейромереж”

Актуальність магістерської дисертатції обумовлена високим розвитком нейромережевих технологій та штучного інтелекту, потреба в зростанні якості прогнозів фінансових ринків, незавершеність формування цілісного уявлення щодо прогнозування цін на фінансових ринках.

Об’єкт дослідження — нейронні мережі на основі рекурентної архітектури, їх можливості та перспективи у сфері фінансового прогнозування.

Предмет дослідження — моделі та методи застосування нейронних мереж для задач прогнозування, шляхи покращення існуючих методів та систем прогнозування.

Метою магістерської дисертації є підвищення якісних характеристик роботи рекурентних нейромереж для розв’язання задач прогнозування фінансових ринків. Для досягнення мети були поставлені наступні задачі:

- Огляд предметної області та аналіз існуючих рішень, архітектур нейромереж;
- Дослідження та порівняння основних концепцій, що пояснюють динаміку цін на фінансових ринках;
- Дослідження науково-методичних засад використання біржової інформації для прогнозування цін на фінансових ринках;
- Дослідження процесу зворотного поширення похибки у часі для простих рекурентних нейромереж;
- Розробка нових підходів до прогнозування на основі використання елементів штучного інтелекту;
- Розробка методів навчання простих рекурентних нейронних мереж, що дозволяють подолати проблему ефекта градієнтного вибуху;
- Розробка програмного комплексу, що забезпечуватиме просте використання існуючих та розроблених методів для вирішення задачі прогнозування часових рядів;

Ключові слова: нейронні мережі, обробка даних, тензор, прогнозування, машинне навчання.

ABSTRACT

Master's thesis: 125 pp., 21 pictures, 25 tables, 8 packs, 34 sources.

The theme of my master's dissertation "Forecasting of financial markets using recurrent neural networks"

The urgency of the Master's dissertation is due to the high development of neural network technologies and artificial intelligence, the need to increase the quality of forecasts of financial markets, the incomplete formation of a holistic view of forecasting prices in financial markets.

Object of research — neural networks based on recurrent architecture, their capabilities and prospects in the field of financial forecasting.

Subject of research — models and methods of using neural networks for forecasting tasks, ways of improving existing methods and forecasting systems.

The purpose of the master's dissertation is to increase the qualitative characteristics of the work of recurrent neural networks for solving tasks of forecasting of financial markets. To achieve the goal were the following tasks:

- Object review and analysis of existing solutions, neural network architectures;
- Research and comparison of key concepts that explain the dynamics of prices in financial markets;
- Research of scientific and methodical principles of the use of stock information for forecasting prices in financial markets;
- Investigation of the process of propagation of error in time for simple recurrent neural networks;
- Development of new approaches to forecasting based on the use of elements of artificial intelligence;
- Development of methods of training simple recurrent neural networks, which allows to overcome the problem of the gradient explosion effect;
- Development of a software complex that will provide easy use of existing and developed methods for solving the problem of prediction of time series;

Key words: neural networks, data processing, tensor, forecasting, machine learning.

ЗМІСТ

ПЕРЕЛІК СКОРОЧЕНЬ.....	8
ВСТУП.....	9
1 ОГЛЯД ІСНУЮЧОГО СТАНУ ПРОБЛЕМНОЇ ОБЛАСТІ.....	11
1.1 СПЕЦИФІКА ПРОГНОЗУВАННЯ ЕКОНОМІЧНИХ ПРОЦЕСІВ.....	11
1.2. ЗАГАЛЬНІ КРОКИ В ПРОЦЕСІ ПРОГНОЗУВАННЯ.....	15
1.3. ЯКІСНІ МЕТОДИ.....	16
1.3.1. Метод ініціативи знизу	16
1.3.2. Дослідження ринку	16
1.3.3. Метод консенсусу.....	17
1.3.4. Історична аналогія.....	17
1.3.5. Метод Делфі.....	17
1.4. МЕТОДИ НА ОСНОВІ ЧАСОВИХ РЯДІВ	18
1.4.1. Наївні методи	21
1.4.2. МЕТОДИ СЕРЕДНЬОГО І ЗГЛАДЖУЮЧІ МЕТОДИ.....	22
1.4.3. Методи екстраполяції тренда.....	25
1.5. Каузальні методи.....	26
1.6. Структурні методи.	27
1.7 ТЕХНІЧНИЙ АНАЛІЗ ФІНАНСОВОГО АКТИВУ	28
1.8 ФУНДАМЕНТАЛЬНИЙ АНАЛІЗ ФІНАНСОВОГО АКТИВУ	29
1.9 АНАЛІЗ ВПЛИВУ НОВИН ТА ЗАПИСІВ У СОЦІАЛЬНИХ МЕРЕЖАХ НА КУРСИ ФІНАНСОВИХ АКТИВІВ	30
1.10 ПРОБЛЕМА ЗНИКАЮЧОГО ГРАДІЄНТУ	31
1.11 ОГЛЯД ПРОГРАМНИХ ЗАСОБІВ, ЩО РЕАЛІЗУЮТЬ НЕЙРОМЕРЕЖЕВІ АЛГОРИТМИ ДЛЯ ВИРІШЕННЯ ЗАДАЧ ПРОГНОЗУВАННЯ.....	33
Висновки до розділу	41
2 МАТЕМАТИЧНІ МЕТОДИ ПРОГНОЗУВАННЯ ФІНАНСОВИХ АКТИВІВ	43
2.1 Визначення тензорних математичних об'єктів, ранг тензора	43
2.2 Тензори другого рангу та їх графічна інтерпретація	45
2.3 Декомпозиція Такера	46
2.4 АНАЛІЗ ТОНАЛЬНОСТІ ПОВІДОМЛЕНЬ ТА НОВИН ПРО ФІНАНСОВІ АКТИВИ	47
2.5 МЕТОДИ ДОСЛІДЖЕННЯ ТОНАЛЬНОСТІ ПОВІДОМЛЕНЬ	49
2.6 Попередня обробка набору даних	50
2.7 Виокремлення ознак.....	51
2.8 Модель N-грамм	51
2.9 Модель “мішок слів”	52
2.10 Методи, що базуються на використанні словників.	53
2.11 РЕКУРЕНТНІ НЕЙРОННІ МЕРЕЖІ.....	54

2.11.1 Різниця між рекурсивними та рекурентними мережами	55
2.11.2 ОБМЕЖЕННЯ РЕКУРЕНТНИХ НЕЙРОННИХ МЕРЕЖ	56
2.12 Довга короткострокова пам'ять (LSTM)	58
2.12.1 Архітектура комірки довгої короткострокової пам'яті.	60
2.12.3 LSTM-архітектури для задачі прогнозування часових рядів	62
Висновки до розділу	64
3 РОЗРОБКА АРХІТЕКТУРИ СИСТЕМИ ПРОГНОЗУВАННЯ.....	66
3.1 Модуль для збору і аналізу даних з соціальних мереж.....	67
3.2 Модуль для збору і аналізу фінансових новин.....	70
4.3 Модуль для збору і аналізу кількісних даних	74
3.4 Модуль для роботи із тензором	76
3.5 НЕЙРОМЕРЕЖЕВИЙ МОДУЛЬ	77
3.5.1 Бібліотеки, що застосовуються у модулі.	77
3.6 Тестування системи	78
4 МОДЕЛЮВАННЯ ТА АНАЛІЗ ОТРИМАНИХ РЕЗУЛЬТАТІВ	87
4.1 Моделювання.....	87
4.2 РЕЗУЛЬТАТАМИ МОДЕЛЮВАННЯ.	87
4.2.1 Оцінка точності прогнозування	88
4.2.2. Оцінка швидкості навчання	90
4.2.3. Оцінка потенційного прибутку від застосування системи	90
Висновки до розділу	95
РОЗРОБКА СТАРТАП-ПРОЕКТУ.....	96
5.1 Опис ідеї проекту	96
5.2 Технологічний аудит проекту	98
5.3 Аналіз ринкових можливостей запуску стартап-проекту	99
5.4 Розроблення ринкової стратегії проекту.....	105
5.5 Розроблення маркетингової програми стартап-проекту	107
Висновки до розділу	111
ВИСНОВКИ ПО МАГІСТЕРСЬКИЙ ДИСЕРТАЦІЇ.....	112
СПИСОК ПОСИЛАНЬ	114

ПЕРЕЛІК СКОРОЧЕНЬ

СУБД — система управління базами даних;

ФЕП — фінансово-економічний процес;

АР — авторегресія;

АРКС — авторегресія з ковзним середнім;

АРІКС — авторегресія з інтегрованим ковзним середнім;

КС — ковзне середнє;

МГУА — метод групового ураховання аргументів;

МНК — метод найменших квадратів;

РМНК — рекурсивний метод найменших квадратів;

СКП — сума квадратів похибок;

СеКП (RMSE) — середньоквадратична похибка;

САП (MAE) — середня абсолютна похибка;

САПП (MAPE) — середня абсолютна похибка в процентах;

АКФ — автокореляційна функція;

ВСТУП

На сьогоднішній день зростає актуальність ефективного вирішення практичних проблем, що включають в себе обробку даних із прихованими кореляціями. У цей клас входить широкий спектр задач, зокрема прогнозування часових рядів, які широко застосовуються у економіці, фізиці та багатьох практичних дисциплінах, неантичний аналіз тексту та ідентифікація природньої мови, нейроаналіз текстовий даних, задачі динамічного регулювання та задачі теорії автоматичного управління.

Варто зауважити, що на сьогоднішній день вже сформовані класичні методи вирішення проблем цього класу у кожній з предметних областей. Наприклад, для прогнозування часових рядів вже більше 30 років застосовують авторегресійний аналіз і метод Бокса-Дженсінса, для регулювання перехідних процесів застосовують ПД-регулятори, а для аналізу тексту — теорію автоматів та граматик. Проте класичні методи часто програють застосуванню нейронних мереж не тільки за швидкістю реалізації, а й за якісними характеристиками, такими як точність прогнозування, час перехідного процесу, точність розпізнавання тексту тощо.

Інтерес наукового суспільства і бізнесу до нейронних мереж зростає з кожним роком. Все більше і більше з'являється прикладів успішного застосування штучного інтелекту у найрізноманітніших галузях діяльності людини, значно зростає кількість підприємств що впроваджують нейронні мережі у свою операційну діяльність. Розв'язання задач управління, прогнозування а також класифікації все частіше покладають на штучний інтелект у багатьох галузях діяльності людини.

Окрім цього, нейронні мережі знайшли широке застосування в прогнозуванні фінансових ринків. Нова і перспективна технологія швидко привернула увагу дослідників. Спочатку нейронні мережі були застосовані у прогнозуванні попиту на певні товари, потім у прогнозуванні ризиків видачі кредитів. Були розроблені системи прийняття фінансових рішень, засновані на штучному інтелекті.

Варто зауважити, що більшість економічних процесів мають нелінійний характер. Нейронні мережі мають можливість працювати не тільки з нелінійними

процесами, а й з зашумленими даними. Стрімкий розвиток нейронних мереж водночас зі стрімким розвитком міжнародних біржових інституцій зумовлює потягу ринку у прогнозуванні фінансових активів.

Актуальність магістерської дисертації обумовлена високим розвитком нейромережових технологій та штучного інтелекту, потребою в збільшенні точності прогнозів фінансових ринків, незавершеністю формування цілісного уявлення щодо прогнозування цін на фінансових ринках.

Метою магістерської дисертації є підвищення якісних характеристик роботи рекурентних нейромереж для розв'язання задач прогнозування фінансових ринків. Для досягнення мети були поставлені наступні задачі:

- Оглянути предметну область та проаналізувати існуючі рішення, архітектури нейромереж;
- Дослідження та порівняння основних концепцій, що пояснюють динаміку цін на фінансових ринках;
- Дослідити науково-методичні засади використання біржової інформації для прогнозування цін на фінансових ринках;
- Дослідити процес зворотного поширення похибки у часі для простих рекурентних нейромереж;
- Розробити нові підходи до прогнозування на основі використання елементів штучного інтелекту;
- Дослідити методи навчання простих рекурентних нейронних мереж, що дозволяють подолати проблему ефекта градієнтного вибуху;
- Розробити програмний комплекс, що забезпечуватиме просте використання існуючих та розроблених методів для вирішення задачі прогнозування тренду часових рядів;

Об'єкт дослідження — рекурентні нейронні мережі, їх можливості та перспективи у сфері прогнозування фінансових ринків.

Предмет дослідження — моделі та методи застосування рекурентних нейронних мереж для задач прогнозування часових рядів, шляхи підвищення точності існуючих методів та систем прогнозування.

1 ОГЛЯД ІСНУЮЧОГО СТАНУ ПРОБЛЕМНОЇ ОБЛАСТІ

1.1 Специфіка прогнозування економічних процесів

Результативність застосування традиційних методів прогнозування фінансових активів, наприклад акцій, облігацій, та валюти, які вільно продаються і купуються на біржах, можна назвати недостатньою для потреб сучасного ринку. Це пов'язано із тим, що інвестиції на фондовому ринку тісно пов'язані із Інтернетом і залежні від інформаційного середовища. Для підвищення точності прогнозування доцільно застосувати таку модель, що не тільки базується на кореляціях факторів та особливостях часового ряду, а й тісно пов'язана з декількома джерелами даних.

Сучасні системи прогнозування не враховують комплексно кількісні та якісні фактори, що впливають на зміну курсів акцій або враховують у векторному вигляді, який обмежує можливості представлення вхідних даних до нейромережі. Такі системи прогнозування зазвичай мають точність 53-58% вірного прогнозу. Цієї точності замало для того, щоб використовувати такі системи як повноцінний інструмент для економічного аналізу та прогнозування. Отже, метою дослідження є підвищення точності прогнозування тренду курсів акцій за допомогою рекурентних нейромереж.

На курси акцій надзвичайно сильно впливають безліч джерел інформації. Нова інформація про компанію чи галузь взагалі може різко змінити настрої інвесторів, що, безумовно, вплине на курс акцій. Взагалі можна сказати, що біржева інформація складна та взаємопов'язана. Дані можуть бути поділені на кількісні та якісні. Кількісні дані (об'єми прибутку, податки, прогноз на наступний квартал тощо) не можуть повністю передати все різноманіття фінансових показників і безпосередньо віддзеркалювати поточний стан речей для компанії та настрої інвесторів. Однак, кількісне прогнозування також стрімко розвивається і покладене в основу цього дослідження. Якісні дані містяться у текстових описах, статтях, повідомленнях, що розміщені у фінансових медіа та соціальних мережах. Якісні дані доповнюють кількісні і формують відносно коректне уявлення про поточний стан речей.

Проблемою залишається дослідження спільного впливу декількох джерел даних на курс акцій. У попередніх дослідженнях наведено два підходи до прогнозування курсів акцій за допомогою рекурентних нейромереж: прогнозування часового без урахування зовнішніх факторів та прогнозування з урахуванням вектору зовнішніх факторів.

1. Прогнозування без урахування зовнішніх факторів являє собою звичайне прогнозування часового ряду за допомогою нейромережі LSTM.

Переваги прогнозування без урахування зовнішніх факторів:

- Швидка реалізація через відсутність модулів аналізу подій, настроїв інвесторів та кількісного аналізу;
- Швидше навчання порівняно з іншими методами;
- Відсутність потреби збирати величезну допоміжну кількість даних. Все що треба для прогнозування — дані часового ряду (курс відкриття, закриття, максимальний, мінімальний, медіана, волатильність);

Недоліки прогнозування без урахування зовнішніх факторів:

- Низька точність прогнозування (зазвичай не більше 54-56%);
- Відсутність гнучкої реакції на фінансові події, що нівелює корись застосування системи у сучасному фінансовому ринку, який керується подіями;

2. Прогнозування із урахуванням вектору зовнішніх факторів поєднує набір факторів у вектор який подається на вхід нейромережі.

Переваги прогнозування із урахуванням вектору зовнішніх факторів:

- Система враховує зовнішні фактори і гнучко реагує на події та фінансові новини;

- Більша точність прогнозування порівняно з попереднім методом (57-59%);

Недоліки прогнозування із урахуванням вектору зовнішніх факторів:

- Вектор входів значно розростається;
- Втрачаються контекстні зв'язки між джерелами інформації;
- Необхідність збирати величезну кількість даних про зовнішні фактори, цінність яких частково втрачається на етапі формування єдиного вектору;

Отже, було проведено багато наукових досліджень на цю тему в останні роки. Загальна ідея минулих досліджень полягала в тому, щоб поєднати набір властивостей декількох джерел інформації у єдиний вектор ознак. Однак це призводило до того що зростали розміри вектора. Окрім того поступово слабнуть або взагалі втрачаються контекстні зв'язки між джерелами інформації.

Інші ж дослідники [2] ігнорували зв'язок між джерелами інформації і спиралися лише на аналіз часового ряду. Таке дослідження давало позитивні результати, однак дослідники що ураховували їх[3] отримували більш точне прогнозування, в середньому на 2-4%. У дослідженні[4] приведений алгоритм, що поділяє зв'язки між джерелами даних на статичні та динамічні. Проте кожне прогнозування у ньому розглядається як окрема задача незалежна від інших прогнозів.

Ціна акцій публічної компанії зумовлена прибутками що має отримати компанія в майбутньому. Неодноразово у наукових дослідженнях [5] було зазначено що фінансові новини дуже сильно впливають на ціну акцій[6]. Деякі дослідники зауважують що фондовий ринок керується подіями[7].

У дослідженні[8] був запропонований метод структуризації події як кортежу, що складається з агенту, предикату і об'єкту а також метод що базується на глибинному навчанні для захвату синтаксичної і семантичної інформації про подію. Існує ціла серія досліджень з застосування результатів аналізу настроїв текстів джерел інформації до волатильності фондового ринку[9]. Методологія аналізу настрою у новинах та соціальних, мережах наведена у дослідженні[10].

У роботі[11] досліджується спільний вплив подій і настроїв інвесторів. Виявлено кореляції між фінансовими подіями та настроями інвесторів. Було доведено доцільність застосування багатовимірних даних для прогнозування фондових ринків. Однак, всі ці дані було поєднано у єдиний вектор. Проте існують скриті зв'язки між джерелами даних що були втрачені у попередніх дослідженнях. Але тензорна структура дозволяє не втрачати ці данні, а навпаки, легко маніпулювати ними і виявляти все більш неочевидні зв'язки.

Перед описом конкретних методів варто розмежувати поняття прогнозування і передбачення. Прогнозування — це процес оцінювання майбутньої події (подій), який в тій чи іншій мірі використовує накопичені попередні дані і об'єднує їх визначеним шляхом, щоб отримати необхідну оцінку. Передбачення, крім цього, оперує суб'єктивними міркуваннями.

Прогнозування неможливе без будь-яких історичних даних. Наприклад, виробник телевізорів на підставі минулих даних може спрогнозувати кількість телевізорів, які необхідно зібрати наступного тижня. Але припустимо, що виробник вирішив налагодити випуск нового товару, для якого у нього немає ніяких даних, для яких він міг би застосувати необхідні моделі і техніки. В даному випадку він не зможе здійснити прогноз, тільки передбачити майбутній результат.

Прогнози часто класифікують по періоду їх дії.

У загальному випадку, короткострокові прогнози (до одного року) використовуються для вирішення поточних операцій; середньострокові (від року до трьох) і довгострокові (понад п'ять років) використовуються в стратегічних цілях.

Варто відзначити, що ідеальний прогноз неможливий через наявність великої кількості факторів, які важко оцінити з високим відсотком точності. Тому, замість пошуку ідеального прогнозу, набагато важливішим є добре знання існуючих моделей і методів та правильне їх застосування в залежності від специфіки даних і предметної області, вміння пристосуватися до неідеального результату прогнозування.

Через те, що прогноз залежить від минулих даних, його надійність і точність будуть знижуватися відповідно до того, як далеко ми хочемо здійснювати прогнозування. Варто відзначити, що точність прогнозу і витрати на його здійснення взаємопов'язані. Кращі прогнози не обов'язково є найточнішими. Такі фактори, як його мета і доступність даних грають важливу роль у визначенні бажаного рівня точності.

Далі, розглянемо класифікацію методів прогнозування, яка представлена на рисунку 1.1.



Рисунок 1.1 — Класифікація методів прогнозування

З рисунку можна зробити висновок, що існує 4 основних класи методів прогнозування, а саме якісні методи, методи на основі часових рядів, казуальні та структурні. Проте, методи містять деякі загальні кроки.

1.2. Загальні кроки в процесі прогнозування

Процес прогнозування в загальному випадку можна розбити на наступні кроки.

- Ідентифікація спільної мети;
- Вибір часового періоду прогнозу;
- Вибір моделі для прогнозу: для цього необхідно володіти знанням про різні моделі, застосування їх в різних ситуаціях, наскільки вони є надійними і в яких даних потребують. Виходячи з цих міркувань, може бути обрана одна або кілька моделей.

- Збір даних: дані повинні бути зібрані і представлені в тому вигляді, якому від них вимагає обрана модель;
- Прогнозування: застосування моделі до зібраних даних і обчислення прогнозу;
- Оцінка: прогноз, отриманий на попередньому етапі, розглядається з урахуванням довірчого інтервалу — більшість моделей дозволяють обчислити верхнє і нижнє значення, між якими розташовується прогнозована величина з певною часткою ймовірності. Також, можуть бути застосовані інші методи для порівняння і поліпшення якості отриманого прогнозу;

1.3. Якісні методи

Якісні методи зазвичай застосовуються при участі групи експертів. Зрозуміло, що система прогнозування фінансових активів не може базуватися на постійній допомозі експертів. Проте, аналіз таких методів може бути корисним для розуміння специфіки задачі прогнозування у економіці[16].

1.3.1. Метод ініціативи знизу

Даний метод будує прогноз послідовним шляхом, починаючи з самого «низу». Спочатку висувається гіпотеза, що людина, що знаходиться «ближче» всього до споживача або кінцевого використання продукту більш компетентна в прогнозуванні майбутнього. Це не завжди так, але в багатьох областях це припущення валідне і використовується в якості базису даного методу. Таким чином, прогнози з нижчого рівня підсумовуються і передаються на рівень вище (наприклад, районне управління). Процедура повторюється до тих пір, поки не досягне найвищого рівня[17].

1.3.2. Дослідження ринку

Фірми часто наймають сторонні компанії, які спеціалізуються у вивченні і дослідженні ринку і займаються даним типом прогнозування. Вивчення ринку найчастіше застосовується для дослідження продукту з точки зору нових ідей, плюсів і мінусів існуючих рішень, аналогічних конкурентних продуктів в даному сегменті і т.д. Найчастіше, методами збору даних в цьому випадку є опитування та інтерв'ю[18].

1.3.3. Метод консенсусу

Ключовою ідеєю даного методу є припущення, що група фахівців, що займають різні позиції, можуть дати більш точний прогноз, ніж маленька група осіб. В даному випадку основою методу є скликання зборів, на яких відбувається вільний обмін ідей і думок осіб з усіх рівнів керівництва («мозковий штурм»). Проблемою даного методу є той факт, що люди, що займають нижчі посади, зазвичай бояться висловлювати думку, відмінну від людей з більш високими посадами[18].

1.3.4. Історична аналогія

Даний метод зазвичай застосовується з урахуванням того, що у нас немає ніяких минулих даних. Наприклад, таке може статися в разі, коли організація готується до випуску нового продукту. Однак у неї можуть бути вже у продажу продукти зі схожими характеристиками. В цьому випадку, можна використовувати накопичені історичні дані для цих продуктів. Обмеження даного підходу полягають у тому, наскільки правдиво припущення про схожість продукту або події, а також змінюються при впливі зовнішніх факторів [18].

1.3.5. Метод Делфі

Даний метод має схожі риси з методом консенсусу, при цьому виправляючи виникаючу при його використанні проблему — упередженість думки. Для її усунення пропонується дотримуватися анонімності осіб, що беруть участь у

опитуванні; в цьому випадку кожен голос має однакову вагу. Кожному учаснику надсилається опитувальник, відповіді якого потім сумуються; учасникам потім надаються результати цього етапу, і висилається новий опитувальник. Зазвичай, даний метод займає близько трьох ітерацій[18]..

1.4. Методи на основі часових рядів

Якщо ми розглядаємо прогнозування фінансових активів, то у більшості ситуацій, що виникають у нас є достатньо накопиченої історичної інформації для прогнозування. Існує безліч методів, які використовують статистичний аналіз використовуючі минулі дані, щоб зробити прогноз на майбутнє. Ключове припущення в даному випадку полягає в тому, що попередні характеристики і взаємозв'язок даних будуть зберігатися і в майбутньому. Різні методи по різному використовують попередні значення і по різному оцінюють їхній внесок в прогнозовані значення.

Значення часового ряду складають фіксовані вимірювання в певні моменти часу конкретної змінної або характеристики. Часові інтервали, між якими відбувається вимір, можуть становити мілісекунди, секунди, хвилини, дні, тижні і т.д.. Наступні визначення дуже важливі для аналізу часових рядів, особливо при прогнозуванні курсів фінансових активів[20]:

1) Тренд: довгострокові зміни в даних, наприклад, зростання ціни або популяції. Приклад наведено на рисунку 1.2:

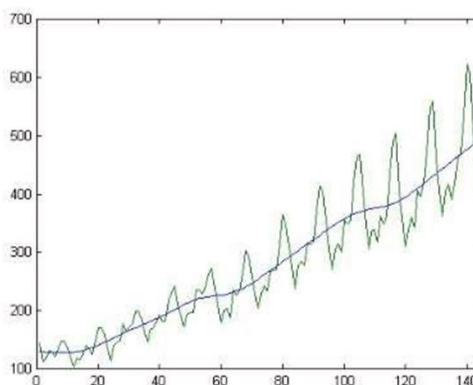


Рисунок 1.2 — Висхідний тренд

2) Сезонні варіації: це можуть бути періодичні варіації в часовому ряді, що виникають із-за так званих сезонних чинників. Наприклад, деякі речі купують більше зимою і менше літом. На ринку акцій сезонні фактори мають особливе значення. Квартальні звіти, розпродажі є саме сезонними факторами. Приклад сезонних змін часового ряду наведено на рисунку 1.3:

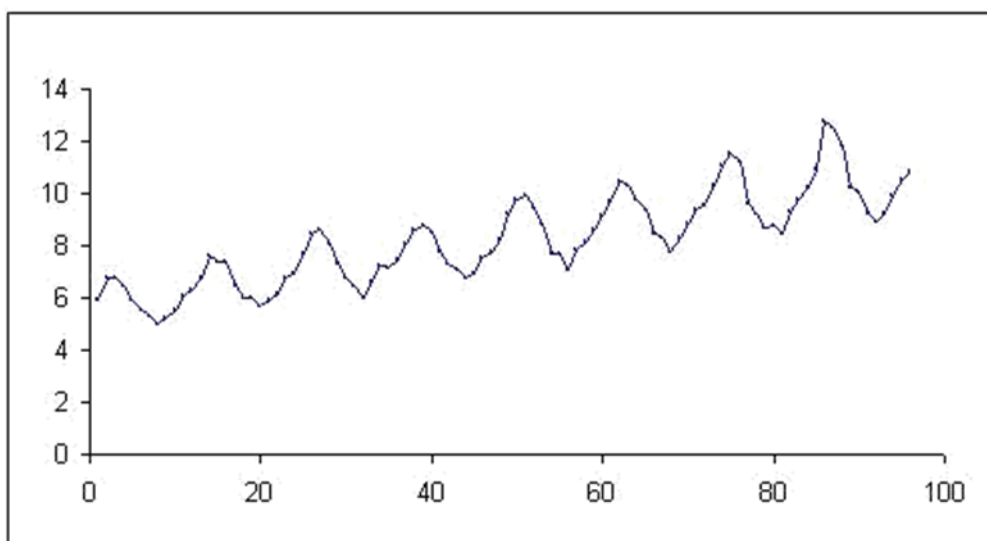


Рисунок 1.3 — Сезонні зміни часового ряду

3) Циклічні варіації: дані варіації виникають через явище під назвою бізнес-цикл. Бізнес-цикл відображає коливання економічної активності даного процесу і може варіюватися від одного року до тридцяти років. Тривалість і рівень змін, викликаних бізнес-циклом зазвичай важко оцінити і врахувати при аналізі часового ряду.

4) Випадкові варіації: зміни в даних, які не можуть бути віднесені до будь-якої категорії з перерахованих. У багатьох випадках, справжня причина цих варіацій може бути знайдена тільки після детального аналізу даних. Подібні зміни можуть бути викликані широким спектром причин, наприклад, зміною погоди, тобто подіями, випадковими по своїй суті. Через це складно передбачити їх вплив на результуючі дані; проте, їх ефект можна усунути, використовуючи згладжуючі техніки. Приклад подібних варіацій представлений на рисунку 1.4:

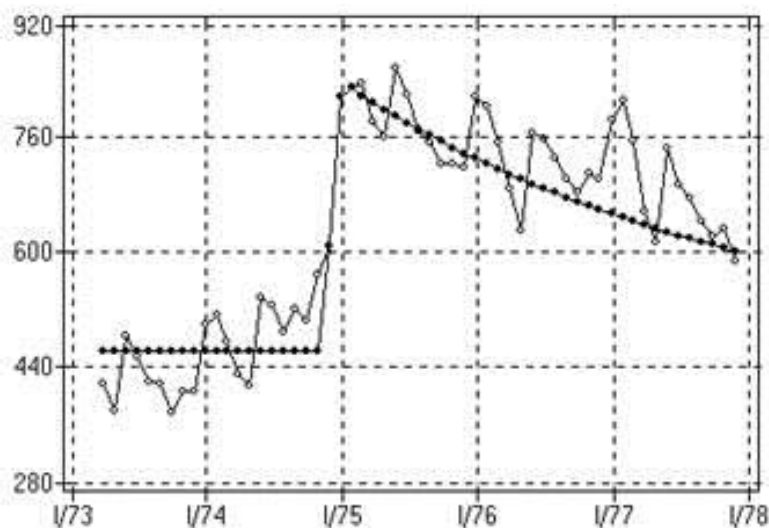


Рисунок 1.4 — Випадкові зміни часового ряду

Часові ряди, як правило, містять всі чотири згадані вище компоненти. Одним з найважливіших завдань в їх аналізі та прогнозуванні є виділення кожної конкретної компоненти з початкових даних з максимальною точністю. Процес виокремлення цих компонентів з часового ряду називається розкладанням часового ряду. Найбільш важливим є виділення тренда шляхом видалення інших компонент. Лінія тренда, як правило, може бути екстрапольована в майбутнє і таким чином може бути отриманий прогноз. Ефект інших компонент на результуючий прогноз може бути врахований при включенні відповідної компоненти.

У більшості короткострокових прогнозів циклічна компонента, як правило, не враховується і її окреме виділення не відбувається. також, передбачається, що випадкові варіації відносно малі і таким чином «гасять» один одного з плином часу. Таким чином, важливою метою в більшості випадків є видалення сезонності з часового ряду. Даний процес називається десезонізацією даних.

Як було сказано вище, існує набір методів, заснованих на аналізі часових рядів і вони по-різному виділяють компоненти з даних. Деякі методи роблять це явно, деякі методи виконують виділення компонентів неявно.

Тепер розглянемо докладніше методи засновані на аналізі часових рядів, та розглянемо їх доцільність до застосування у системі прогнозування фінансових ринків.

1.4.1. Наївні методи

Методи, включені в цю групу, з точки зору математичного опису є найпростішими методами прогнозування. Найпростіший полягає в тому, що ми просто беремо останнє виміряне значення і використовуємо його в якості оцінки прогнозу. При цьому всі попередні значення часового ряду ніяк не враховуються.

Для даних з високою сезонністю може бути використаний схожий метод, тільки в цьому випадку, прогноз наступного значення ряду приймається за останнє значення ряду в такому ж сезоні (наприклад, для всіх наступних прогнозів на серпень беруться значення за останній серпень) [21]:.

Також, варіацією першого методу є так званий метод дрейфу, який дозволяє оцінці прогнозу збільшуватися або зменшуватися з плином часу; зміни (так званий дрейф) в даному випадку рівні середньому зміні в попередніх даних, що схоже з проведенням прямої лінії між першим і останнім вимірами і екстрапольоване далі.

Також, одним з найпростіших методів прогнозування є так званий метод середнього, коли за прогнозоване значення береться середнє по всім попереднім даним.

Приклади роботи даних методів наведені на рисунку 1.5.

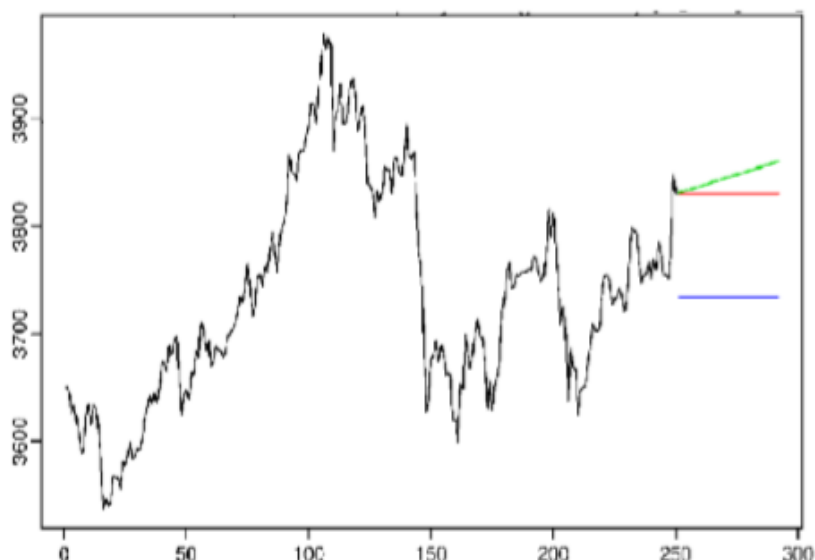


Рисунок 1.5 — Часовий ряд і його прогнозовані значення. Синім позначено метод середнього; червоним — простий наївний метод; зеленим — метод дрейфу.

1.4.2. Методи середнього і згладжуючі методи.

Коли часовий ряд не змінюється занадто швидко і не має сезонних характеристик, то метод змінного середнього може бути використаний з метою видалення випадкових варіацій і здійснення оцінки прогнозу. У цьому випадку, ми використовуємо величину так званого змінного (локального) середнього. Її можна вважати відносним компромісом між простою моделлю середнього і методом дрейфу.

В даному випадку дуже важливим є правильний вибір періоду для оцінки змінного середнього, так як, наприклад, при виборі великого часового періоду і наявності невеликого тренда дані будуть «відставати» від тренда, в той час як занадто короткий період буде надавати більшого значення різним варіаціям[20]:.

Формула для моделі простого змінного середнього виглядає наступним чином:

$$Y_{t-1} = \frac{Y_t + Y_{t-1} + Y_{t-2} + \dots + Y_{t-m+1}}{m} \quad (1.1)$$

де Y_t — значення моделі;

m — кількість ітерацій.

Можна відзначити наступні важливі характеристики даного методу:

- Різні ковзаючі середні будуть давати нам різну оцінку прогнозу;
- Чим вище число спостережень, які використовуються для даної моделі,

тим вище так званий ефект згладжування;

- Якщо в даних спостерігається тренд, який можна вважати постійним з невеликими випадковими варіаціями, то тоді, як правило, можна використовувати більшу кількість спостережень;

- В іншому випадку, якщо є припущення, що в даних в майбутньому може статися значна зміна, то тоді нам знадобиться в більшій мірі розраховувати на найближчі значення і, як правило, варто обмежитися меншим числом спостережень;

Обмеження даного підходу полягають у наступному:

- В даному випадку кожному значенню дається рівна вага і воно вносить рівний внесок в кінцеву оцінку прогнозу, в той час як логічно припустити, що більш недавні значення є більш значущими в даній ситуації;

- Даний метод не бере до уваги дані, що знаходяться поза часового проміжку, використовуваного для підрахунку змінного середнього, що означає, що не всі доступні нам дані використовуються в повній мірі;

- У разі наявності сезонних варіацій прогноз може показувати не дуже хороший результат;

Описаний вище підхід можна змінити, надавши кожному значенню певну вагу. У цьому випадку така модель буде називатися моделлю зваженого ковзного середнього. Її формула приведена далі:

$$F_t = w * Y_t + w * Y_{t-1} + w * Y_{t-2} + \dots + w * Y_{t-m+1} \quad (1.2)$$

В даному випадку більшість значень можуть бути проігноровані (отримавши нульову вагу), а конкретні ваги можуть бути обрані, виходячи з різних припущень

(наприклад, більш старі дані можуть мати більш високу вагу). Сума всіх ваг повинна дорівнювати одиниці.

У більшості випадків для вибору ваг керуються правилом, що недавні значення є більш важливими і вносять більший внесок в оцінку прогнозу, тому для них беруться великі ваги. Однак, якщо ми маємо справу з сезонністю в даних, то це слід враховувати і брати до уваги.

Попередні два методи мають один спільний недолік — ми повинні зберігати велику кількість історичних даних. Додаючи нові значення в ці методи, ми тим самим відкидаємо старі значення. У більшості додатків, як було згадано вище, прийнято вважати, що недавні значення є більш інформативними і несуть більше корисної інформації, ніж старіші. Якщо це припущення є вірним, то метод експоненціального згладжування є найбільш логічним і простим методом в цій ситуації[21]:.

Даний метод автоматично привласнює ваги попереднім значенням таким чином, що ваги зменшуються експоненціально. Основний принцип даного методу полягає в тому, що ми вважаємо, що оцінка прогнозу наступного значення складається з двох компонент: старої оцінки прогнозу і певної величини помилки прогнозу, що можна виразити наступною формулою:

$$Y_t = \alpha * Y_t + w * Y_{t-1}(1 - \alpha) \quad (1.2)$$

Також наведемо формулу обліку декількох попередніх значень:

$$Y_t = \alpha[Y_t^2 + \dots + w * Y_{t-1}(1 - \alpha)^3] \quad (1.3)$$

Величина константи згладжування може варіюватися між нулем і одиницею. Чим вище його значення, тим більш оцінки прогнозу стають більш чутливими до поточних умов і навпаки, при маленькій величині прогноз більш стабільний і менше реагує на поточний стан.

На поточний момент даний метод є найбільш часто використовуваним методом прогнозування і включений в більшість програм здійснення прогнозу. Наступні фактори забезпечили цим методом таке успішне застосування:

- даний метод є досить точним
- формула, що описує основну модель, відносно проста
- для його застосування потрібно відносно мало обчислювальної потужності

Також варто відзначити, що у даного методу є різні розширення, які допомагають враховувати такі речі як сезонність і наявність тренда. Вони включають в себе подвійне або потрійне експоненціальне згладжування, а також, наприклад, метод Холта, який враховує наявність тренда.

1.4.3. Методи екстраполяції тренда

Дані методи прогнозування намагаються вписати в існуючі дані лінію тренда і потім екстраполювати її для майбутніх значень прогнозу. Існують різні види рівнянь тренда: лінійні, експоненціальні, квадратичні і т.д. З усіх складових часового ряду компонент саме тренд задає довгостроковий напрямок.

Розглянемо поширений метод прогнозування з даної категорії — метод простої лінійної регресії. Регресія в загальному випадку може бути описана як функціональний взаємозв'язок між двома і більше корелюючими змінними. Іншими словами, прогноз однієї змінної відбувається, виходячи зі значень іншої. Дане твердження можна виразити в такій формулі:

$$Y_t = b_0 * Y_t + b_1 * Y_{t-1} \quad (1.4)$$

Лінійна регресія найчастіше застосовується для довгострокового прогнозування.

Лінійна регресія застосовується як метод прогнозування на основі часових рядів, так само як і каузальний метод прогнозування. Це залежить від того, з якою

залежною змінною ми маємо справу. Якщо вона змінюється в залежності від часу, то ми маємо справу з часовими рядами; якщо ж змінна змінюється в залежності від іншої змінної, то тут йде мова про каузальний взаємозв'язок. Приклад застосування методу лінійної регресії наведено на рисунку 1.7:

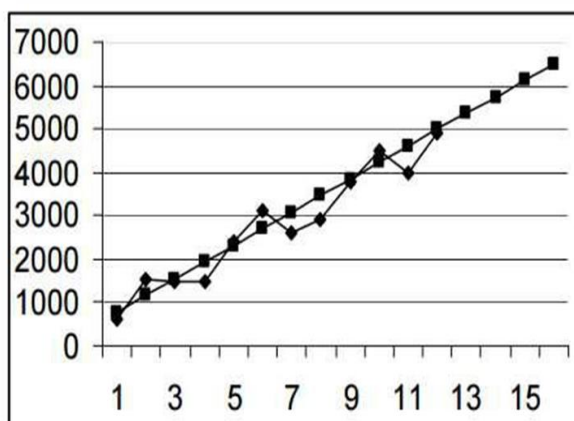


Рисунок 1.7 — Застосування методу лінійної регресії

Даний підхід не бере до уваги такі фактори, як сезонність і циклічність даних.

Найчастіше, для оцінки коефіцієнтів даного методу використовується так званий метод найменших квадратів. Суть методу полягає в тому, що необхідні значення підбираються таким чином, щоб мінімізувати суму квадратів різниць між конкретним спостереженням і відповідної їй точки на лінії.

1.5. Каузальні методи

Дані методи намагаються знайти та ідентифікувати фактори, які можуть вплинути на оцінку прогнозу даної змінної. Наприклад, якщо взяти до уваги інформацію про характеристику клімату, то теоретично можна краще спрогнозувати попит на парасольки.

Подібні методи припускають, що подібні фактори можуть бути виміряні та враховані при оцінці прогнозу. Мета подібних методів — виробити найкращий статистичний взаємозв'язок між розглянутою залежною змінною і однією або декількома незалежними.

Найпоширеніший метод в даній категорії є метод регресійного аналізу, який ділиться на простий (залежність від однієї змінної) і багатовимірний (залежність від декількох змінних).

Подібні методи найчастіше застосовуються для середньострокового і довгострокового прогнозування.

Крім описаних вище, варто згадати метод комп'ютерної імітації, який, як правило, застосовується на великих обсягах даних[21]:.

1.6. Структурні методи.

У даних методах визначальною ідеєю є припущення, що в основі спостережуваного процесу лежить якась структура, для якої задані деякі параметри, і, що найважливіше, базові правила і залежності.

До подібних методів можна віднести наступні:

- моделі на основі нейромереж;
- моделі, що базуються на використанні Марковських ланцюгів;
- моделі, що використовують класифікаційно-регресійні дерева;
- моделі, засновані на використанні генетичних алгоритмів;
- моделі, що використовують опорні вектора;
- моделі, побудовані із застосуванням апарату нечіткої логіки.

Саме ці методи є найбільш придатними до застосування у системах прогнозування фінансових активів. Нейромережеві методи набирають все більшої популярності для вирішення задач прогнозування у останні роки. Дослідження[] показало, що застосування LSTM-нейромереж до часового ряду може давати прогнозування тренду вірно у 58% випадках. Саме тому у цій магістерській дисертації обрано нейромережі як засіб прогнозування. Проте при прогнозуванні фінансових активів часто застосовують також технічний і фундаментальний аналіз.

Технічний і фундаментальний аналіз є невід'ємними складовими аналітики фінансових ринків. Дослідники фінансових ринків використовують

фундаментальний та технічний аналіз для прогнозування основних фінансових показників, курсів акцій а валют. Ці методи є класичними і традиційними вирішеннями задачі прогнозування фінансових ринків і успішно застосовуються вже більше тридцяти років.

1.7 Технічний аналіз фінансового активу

Метою технічного аналізу є дослідження закономірностей коливань минулих і теперішніх технічних показників змін певного активу для подальшого прогнозування коливань цих показників. В основі технічного аналізу покладена гіпотеза про те, що коливання цін акцій часто містить певні закономірності, які у технічному аналізі називаються тенденціями руху фінансового активу. При цьому під час встановлення певної тенденції для активу доцільно виконувати характерні для окремої тенденції ринкові операції та використовувати характерні для кожної тенденції стратегії торгівлі. Наприклад, якщо фінансовий актив має спадаючий тренд протягом деяких тактів технічного аналізу, системи технічного аналізу зазвичай рядять здійснити продаж активу із його подальшою повторною купівлею через 2-3 такти аналізу, або коли відбудеться зміна тренду[22]:. Таким чином, технічний аналіз є інструментом прийняття рішень, що спирається лише на дані графіку і технічні дані фінансового активу, не беручі до уваги такі важливі чинники, як фінансові новини, інфоповоди, повідомлення людей у соціальних мережах. Частково цей недолік перекидає загальна доступність інструментів технічного аналізу. Більшість систем технічного аналізу не лише безкоштовні і загальнодоступні, а й інтегровані у більшість сучасних біржових ресурсів. Також технічний аналіз дає змогу дізнатися, чому саме система радить прийняти те чи інше рішення. За даними соціологічних опитувань, більшість досвідчених трейдерів не спирається при прийнятті рішень на дані технічного аналізу через низьку точність результатів, застарілість та загальнодоступність технології. Проте інструменти технічного аналізу можуть бути корисні для дослідників економічних процесів,

зокрема як допоміжний засіб при прогнозуванні фінансових активів за допомогою рекурентних нейронних мереж.

1.8 Фундаментальний аналіз фінансового активу

Фундаментальний аналіз розглядає вирішення задачі прогнозування з іншого боку. На відміну від технічного, він спирається не на закономірності у графіках та послідовностях, а на дослідження подій, що стосуються фінансового активу, політичних та економічних новин. Центральний об'єкт дослідження фундаментального аналізу — компанія із якою пов'язаний ресурс та середовище у якому вона перебуває. Фундаментальний аналіз враховує ризики, прогнози продажів, фінансові звіти. Зазвичай, фундаментальний аналіз проводять за наступним алгоритмом. Спочатку досліджують економічну ситуацію в світі, потім в країні де представлена компанія-емітет фінансового активу[22]. Після цього досліджують ситуації в галузі де веде діяльність обрана компанія, краї які задіяні у виробничих процесах. Потім досліджуються фінансові показники компанії за деякий період часу (зазвичай квартал) та компаній-партнерів та постачальників обраної компанії. В кінці досліджуються прогнози продажів продукту компанії в майбутньому і визначається цільова вартість акції компанії. Так, наприклад, при фундаментальному аналізі акцій компанії Apple спочатку буде досліджена економічна ситуація в світі, потім у США. Потім буде досліджена економічна і політична ситуація у країнах які беруть участь у виробничому процесі продукту компанії (Китай, Індія). Потім буде досліджена ситуація в галузі інформаційних технологій, зокрема ринки виробництва смартфонів, комп'ютерів, планшетів та переносної електроніки. Після цього будуть досліджені квартальні показники компанії та компаній-постачальників (наприклад, Samsung) та прогнози продажів і прибутку.

Із сказаного вище можна зробити висновок, що під час проведення фундаментального аналізу багаторазово виконується операція кількісної оцінки деякої інформації. З цього випливає нова задача — кількісна оцінка тональності

текстової інформації і економічних показників відносно фундаментального аналізу. Зрозуміло, що це може виконати людина-експерт. Проте з появою нейронних мереж і нейромережових засобів аналізу тексту ситуація змінилася. Отже, однією з проблем що розглядається у магістерській дисертації є формалізація і оцінка тональності основних факторів фундаментального аналізу для покращення точності прогнозування фінансових активів за допомогою рекурентних нейронних мереж.

Проведення фінансового аналізу спрощують десятки економічних та політичних індексів, що оцінюють економічну ситуацію в світі на певній країні, ймовірність терористичних загроз, війни, ставки рефінансування тощо. Проте, цей метод залишається одним з найважчих у реалізації але і одним із найточніших. Ускладнює аналіз те що різні чиники у різний період часу можуть мати різний вплив на результат. Проте, застосування нейромереж, а також застосування у парі із технічним аналізом може нівелювати цей недолік.

1.9 Аналіз впливу новин та записів у соціальних мережах на курси фінансових активів

За даними соціологічних досліджень, 80% інвесторів-початківців та 53% інвесторів із досвідом більше 10 років при виборі активу для інвестицій у якості одного з основних факторів розглядають, серед найважливіших, реакцію користувачів соціальних мереж та повідомлення, що стосуються певного активу.

Деякі інвестори та трейдери надають переваги стратегії “слідкування за професіоналом”. Сутність цієї стратегії полягає в тому, що інвестор або трейдер намагається повторити дії більш досвідченого та успішного, часто всесвітньо відомого, учасника економічного ринку з метою отримання такого самого прибутку у відсотках як і у “оригінала”. Безумовно, така стратегія має низку вагомих недоліків. По-перше саме учасник ‘оригінал’ часто буде отримувати значно більший прибуток у відсотках, ніж учасник торгівлі, що повторює його дії. Це пов’язано з том, що інформація про дії відомого торговця провокує масові дії з боку торговців, що слідкують за ним, і, відповідно значне зростання або спадання курсу вже після

того як ‘оригінал’ купив чи продав актив відповідно. Оскільки торгівці що слідкують повторюють дію з деякою затримкою від оригінала, вони купують чи продають актив за значно менш вигідною ціною і отримують значно менший прибуток (часто збиток). Так новина про те, що Уорен Баффет розширює свій пакет акцій компанії Apple на 50% у квітні 2018 року збільшила кількість угод з продажу акцій у той день на 3 млн, а сам курс акцій виріс на 2,57%. Новина про те, що Уорен Баффет продає всі свої акції компанії IBM у жовтні 2017 року спровокувала 700 тисяч додаткових угод і обвалила курс акції компанії на 4,56%. Незважаючи на не оптимальність цієї стратегії, ці данні є надцінними для нашої системи через масове застосування цієї стратегії учасниками ринку.

Загалом, можна зауважити, що рішення більшості людей про купівлю чи продаж певного активу спирається на інформацію яку вони отримують з мережі Інтернет. У більшості випадків ця інформація є публічною, тобто всі мають до неї необмежений доступ.

Серед основних критеріїв які доцільно застосувати при аналізі повідомлень у соціальних мережах доцільно виділити два — кількість і тональність. При тому кількість, а точніше збільшення кількості є первинним критерієм, а тональність — вторинним. Зрозуміло що тональність це перш за все позитивна чи негативна думка автора, у той час як кількість показує загальну зацікавленість ринку у певному активі.

Для аналізу тональності доцільно ввести таку метрику, яка буде враховувати позитивні та негативні словосполучення у повідомленнях на обчислювати різницю між позитивними та негативними. Такі списки вже творені для англійської мови і знаходяться у стадії активної розробки для російської та української. У дослідженні такий список було взято з відкритих джерел.

1.10 Проблема зникаючого градієнту

Проблема зникаючого градієнта це складність, яка виникає при навчанні штучних нейронних мереж з використанням методів навчання на основі градієнта і

зворотного поширення похибки. Стандартні функції активації, такі як гіперболічний тангенс, мають градієнти в діапазоні $(-1, 1)$, а метод зворотного поширення помилки обчислює їх за ланцюговим правилом. Після множення цих чисел для обчислення градієнтів "фронтальних" шарів в n -шаровій мережі, градієнт (сигнал помилки) експоненційно зменшується разом з n , а передні шари навчаються дуже повільно[23].

Коли використовуються функції активації, похідні яких можуть приймати великі значення, є ризик зіткнутися з проблемою вибухаючого градієнту. Можливими рішеннями є:

Багаторівнева ієрархія: шар мережі попередньо навчається, використовуючи методи навчання без учителя, а потім його значення регулюється за допомогою методу зворотного поширення помилки.

Довга короткострокова пам'ять: різновид архітектури рекуррентних нейронних мереж. Коли величини помилки поширюються в зворотному напрямку від вихідного шару, помилка не випускається з пам'яті LSTM-блоку. Вона безперервно передається назад кожному з вузлів, поки вони не будуть навчені відкидати подібні значення;

Залишкові мережі (Residual networks): один з найбільш ефективних методів вирішення проблеми зникаючого градієнта є використання залишкових нейронних мереж (ResNets). Більш глибока мережа матиме вищу помилку навчання, ніж мережа з меншою кількістю шарів. Команда Microsoft Research виявила[], що поділ глибокої мережі на частини (скажімо, кожна частина являє собою три шари мережі) і передача вхідних даних в кожен фрагмент до наступного фрагмента (поряд із залишковим виходом) допомогли усунути більшу частину проблеми зникнення градієнта. Ніяких додаткових параметрів або змін в алгоритмі навчання не потрібно. ResNets показали нижчу помилку навчання (і тестову помилку), ніж їх менші за кількістю шарів аналоги.

Дослідивши методи боротьби із проблемою зникаючого градієнту, можна зробити висновок, що одне з можливих рішень — а саме варіація рекуррентних нейромреж LSTM позбавлена цієї проблеми. Це підтверджує доцільність вибору

мереж із довгою короткострокову пам'яттю у якості інструменту прогнозування фінансових активів.

1.11 Огляд програмних засобів, що реалізують нейромережеві алгоритми для вирішення задач прогнозування

На сьогоднішній день вже розроблено велику кількість програмних продуктів, придатних для застосування там, де виникає необхідність використання нейромережевих технологій у прогнозуванні фінансових ринків. Існують універсальні нейромережеві пакети, призначені для вирішення будь-яких завдань прогнозування (Brain Maker Pro, NeuroSolution), але, як показує практика, такі програмні продукти не завжди зручні для вирішення задач прогнозування часових рядів і курсів фінансових активів.

Існує окремий клас нейромережевих програмних продуктів, призначених виключно для вирішення завдань прогнозування курсів акцій. Ці продукти орієнтовані на фінансових працівників та людей, зацікавлених у дослідженні фінансової галузі — трейдерів, біржових аналітиків і інших. Часто такі програмні пакети мають дружній графічний інтерфейс, і проектуються таким чином, щоб людина, що має слабе поверхове уявлення про нейронні мережі, змогла швидко їх освоїти. До таких програмних продуктів належать: Neuro Builder 2015, NeuroShell Day Trader, BioComp Profit, NeuroScalp.

Найбільш популярні у світі програмні продукти для вирішення задач прогнозування часових рядів, що реалізують нейромережеві підходи: Brain Maker Professional, NeuroShell Day Trader, Neuro Builder 2015.

Пакет Brain Maker Professional — призначений для побудови нейронних мереж зворотного поширення. Пакет включає в себе програму підготовки та аналізу вихідних даних NetMaker, програму побудови, навчання і запуску нейромереж BrainMaker, а також набір утиліт широкого призначення. Програмний пакет орієнтований на широке коло завдань — від створення прогностичних застосувань до організації систем розпізнавання образів і нейромережевої пам'яті. Значна кількість

функцій програми орієнтована на фахівців в області дослідження нейромереж. Слід зазначити, що організація внутрішніх нейромережових моделей є "прозорою" і легко доступною для програмного розширення. У програмі BrainMaker передбачена система команд для пакетного запуску. Існує інтерфейсна програма-функція для включення навчених мереж в програми користувача. В цілому пакет може бути інтегрований в програмний комплекс цільового використання.

Програма BrainMaker призначена для побудови нейромережі за деякими початковими умовами, її навчання в різних режимах, модифікації параметрів мережі. Програма має значну кількість функцій для оптимізації процесу навчання. Крім цього, програма надає ряд методів аналізу чутливості виходів мережі до різних варіацій вхідних даних, при цьому формується детальний звіт, згідно з яким можна додатково оцінити ступінь функціональної залежності вхідних і вихідних значень.

NeuroShell Day Trader — нейромережева система, яка враховує специфічні потреби трейдерів і досить легка у використанні. NeuroShell Trader має, як і в інші стандартні програми для трейдерів, дружній графічний інтерфейс користувача. Можливості графіки дозволяють відображати дані у вигляді японських свічок (candlestick), в формі open / high / low / close, high / low / close, лінійних графіків або гістограм різних типів. Існує можливість міняти кольори, стискати і розтягувати шкали, ховати і знову робити видимими потоки даних. У щоденній, тижневої або місячної часовій шкалі можна відображати курси акцій, ціни товарів (commodities), біржові індекси (indexes), взаємні фонди (mutual funds), обмінний курс валют (foreign exchange rates) і подібне. Інтерфес користувача програми зображено на рисунку 1.8

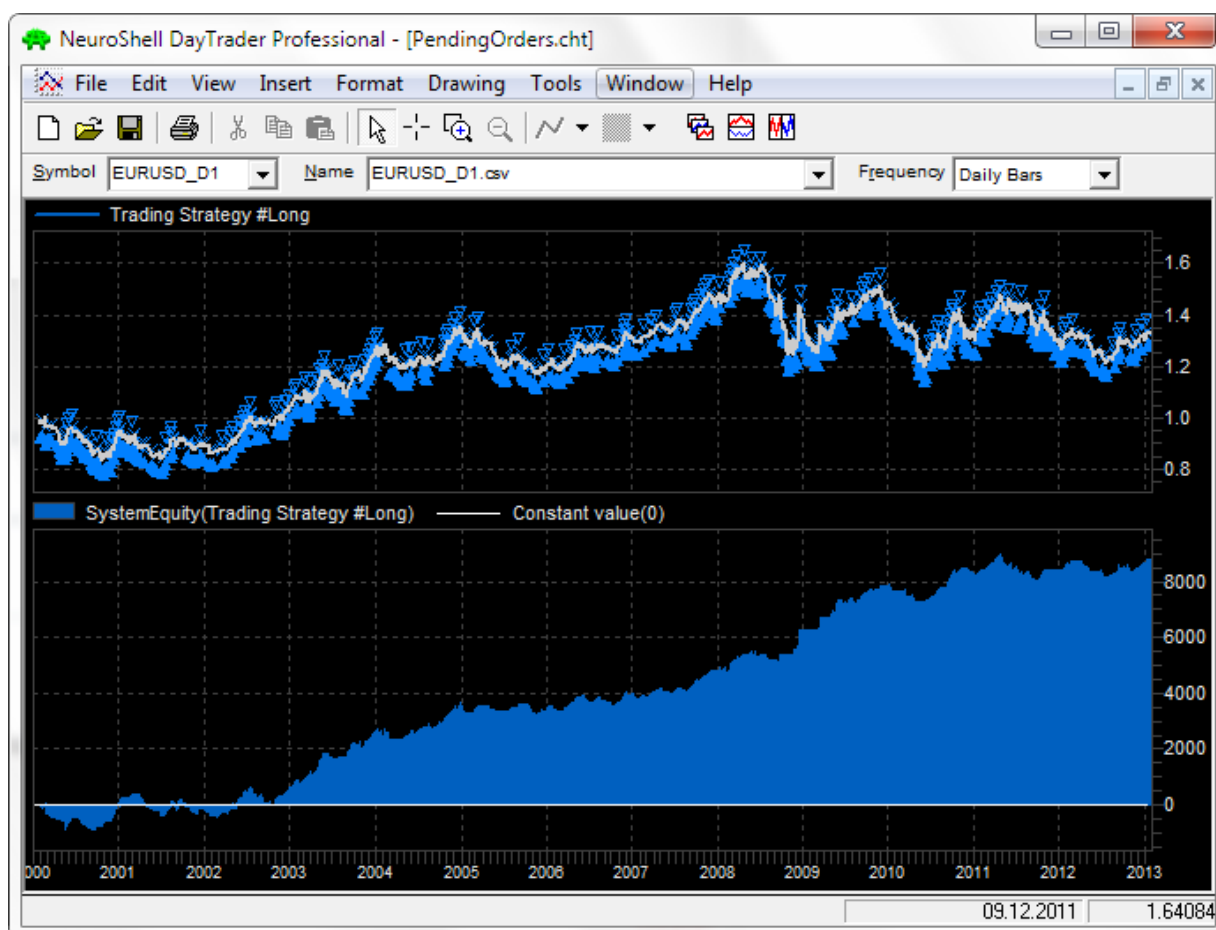


Рисунок 1.8 — Графічний інтерфейс NeuroShell DayTraider

NeuroShell Trader з легкістю читає стандартні текстові файли формату open / high / low / close / volume, які постачають більшість відповідних агентств. Зокрема, NeuroShell Trader працює з текстовими файлами в форматі MetaStock (включаючи версію 6) і файлами даних в форматі, використовуваному програмами TradeStation, SuperCharts і Wall Street Analyst, які Omega Research поширює через власні канали збуту. Найчастіше ці дані безпосередньо можуть бути використані в якості вхідних змінних для нейронної мережі.

У NeuroShell Trader є велика бібліотека з більш ніж 800 технічних індикаторів. Крім стандартних індикаторів, таких як ковзаючі середні (moving averages), норма зміни (rate-of-change) або стохастичні лінії (stochastics), NeuroShell Trader дає можливість реалізувати власні індикатори шляхом комбінації готових функцій зі значного списку, до якого входять умови «якщо -то », арифметичні оператори, тригонометричні функції і багато іншого.

Однією з головних переваг даного продукту є те, що нейронні мережі є вбудованими, а не є чимось привнесеним ззовні чи використовуваним окремо. Вони присутні в меню під рубрикою "Predictions" (Прогнози) поряд з "Indicators" (Індикаторами) і "Data" (Даними). Майстер прогнозу (Prediction Wizard) дозволяє вибрати, що користувач хоче прогнозувати. Це можуть бути ціни закриття (close), їх відсоткові зміни чи інші дані або індикатори. Існує можливість встановлювати, на скільки днів вперед будувати прогнози.

Neuro Builder 2015 Advanced — продукт, що належить до категорії наукомістких, високотехнологічних, вузькопрофесійних інструментів. Це додаток, що працює під управлінням ОС Windows. У своїй галузі — спеціалізовані програми для фінансових аналітиків — Neuro Builder 2015 займають прикордонне положення між серійними програмами і системами на замовлення. Вона може бути використана як самостійний продукт, може виступати складовою частиною складного аналітичного комплексу. Нижче перераховані сім головних характеристик програми Neuro Builder 2015, що відрізняють її від аналогів та конкурентів:

а) програма Neuro Builder 2015 — додаток, створений спеціально для вирішення завдань прогнозування на фінансових ринках;

б) програма Neuro Builder 2015 — додаток, що дозволяє використовувати нейромережі в повсякденній роботі так само просто, як і звичні для трейдерів інструменти — програми технічного аналізу та електронні таблиці;

в) програма Neuro Builder 2015 — додаток, що дозволяє користувачеві використовувати навички, набуті під час роботи з Microsoft Office — технологічність і регулярність, які забезпечуються автоматизацією роботи програми за звичними для користувача сценаріями і оформлення звітів програмою за створеними ним шаблонами;

г) програма Neuro Builder 2015 — єдиний на сьогоднішній день серійний та оновлюваний програмний продукт, що містить модуль дослідження даних до визначення архітектури нейромережі — Best Builder, що дозволяє автоматизувати визначення вхідного вектора параметрів завдання з урахуванням впливу кожного параметра вхідного вектора на передбачуваний результат;

д) програма Neuro Builder 2015 — не є «чорним ящиком»; Детальна документація містить опис всіх структур і файлів, включаючи тимчасові файли. Всі файли з даними системи зберігаються тільки в двох форматах — текстовий і EXCEL;

е) програма Neuro Builder 2015 — продукт, що не залежить від джерела даних; до складу програми входить модуль Data Builder Light, що дозволяє перетворювати фінансові дані з безлічі популярних форматів, в формат даних програми Neuro Builder 2015 і виправляти помилки в даних паралельно з їх перетворенням;

ж) програма Neuro Builder 2015 — дозволяє використовувати знайдені рішення неодноразово; дані для конкретного завдання завжди формуються на етапі її рішення через запит до бази даних і відразу знищуються після отримання результату; між сеансами роботи зберігається тільки опис способу отримання даних з локальної бази;

з) програма Neuro Builder 2015 реалізована у вигляді безлічі незалежних модулів, взаємодіючих в рамках комплексу за документованими інтерфейсами; кожен з програмних модулів оформлений у вигляді виконуваної програми (EXE) і відповідає за вирішення однієї з конкретних підзадач в складі загальної задачі прогнозування на фінансових ринках; кожен з модулів в змозі працювати як в складі комплексу програми Neuro Builder 2015, так і спільно з будь-якими іншими програмами, що підтримують її інтерфейс.

Технологія застосування програми Neuro Builder 2015 орієнтована на регулярність отримання результатів і економію робочого часу аналітика. Так програма Neuro Builder 2015 забезпечує мінімальний період прогнозування, відповідний однієї доби. В кінці торгового дня в базу даних програми заносяться ціни поточного дня, і скрипт буде працювати на обробку нових даних за заздалегідь підготовленим сценарієм. Контроль оператора в процесі обрахунків не потрібен. На початку наступного торговельного дня за результатами обрахунку вже можна отримати прогноз цін закриття цього дня. Таким чином, основний час роботи програми припадає на ніч, і завдання — прогнозування на день вперед — вирішене. Роль користувача полягає в підготовці коректних сценаріїв для роботи програми і

забезпеченні безперебійної подачі живлення персонального електронної обчислювальної машини, на якій запущена програма. Графічний інтерфейс користувача програми зображено на рисунку 1.9

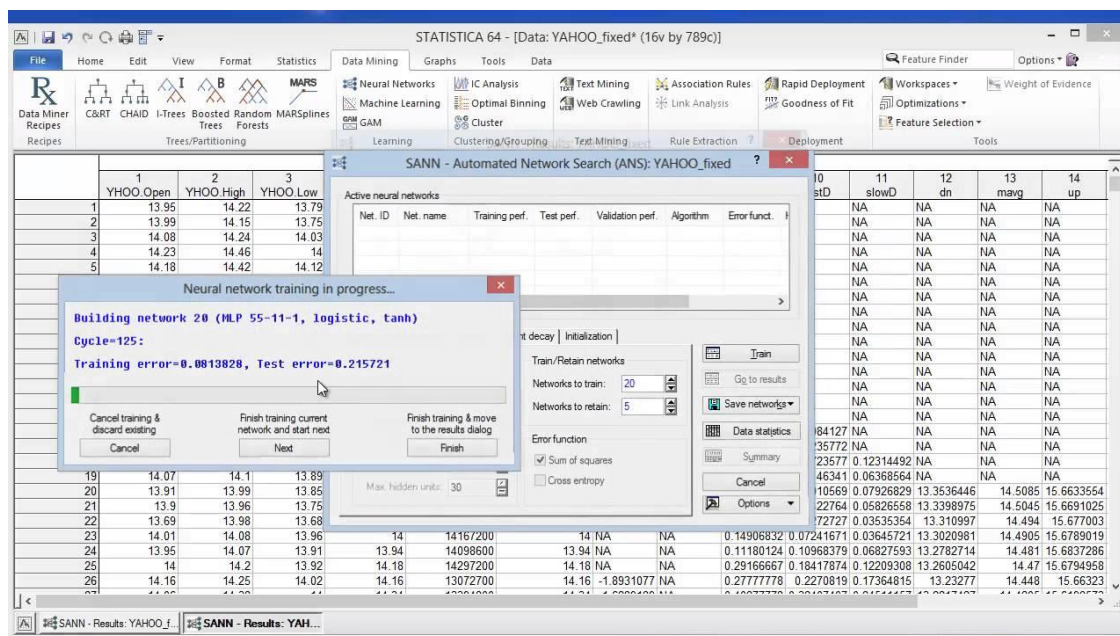


Рисунок 1.9 — графічний інтерфейс Neural Builder

Інші програмні продукти, що реалізують нейромережеві алгоритми, менш відомі і не набули широкого поширення. Таким чином, можна зробити висновок, що типовий програмний продукт ринку нейромережевих програм, призначений виключно для прогнозування фінансових ринків, оцінюється виробниками приблизно в \$800 — \$3000. Зазначена ціна за для українських споживачів досить висока, проте, порівняно з потенційними прибутками, які можна отримати із використанням розглянутих програмних продуктів при успішній торгівлі на фінансових ринках, це дуже незначна сума.

Важко оцінити наведені вище системи за одним критерієм — точність прогнозування, оскільки цей показник часто подається у завищеному вигляді через його маркетинговий вплив. Також майже неможливо об'єктивно співставити потенційний прибуток із ціною ліцензії, оскільки це залежить від великої низки зовнішніх факторів, в основному — від кваліфікації оператора системи. Також, варто зауважити що описані програмні комплекси орієнтовані на різних

користувачів. Порівняння описаних вище застосувань із розробленою системою наведено у Таблиці 2.1

Таблиця 2.1 — Порівняння аналогів із розробленою системою

	Система прогнозування				
	Neuro Builder	NeurShell trTrader	BioComp Profit	NeuroScalp	Розроблена система
Вартість	\$3500	\$800	\$2500	\$850	-
Рік релізу	2015	2013	2010	2015	2018
Вбудована обробка зовнішніх джерел	+	+	-	+	+
Універсальність (можливість застосування до будь якого ринку)	+	+	+	+	+
Гнучкість (можливість індивідуального налаштування)	+	+	-	+	-
Інтеграція з популярними програмними продуктами	-	Excel, LibreOffice	Excel	Excel	-

Продовження таблиці 2.1

Використання постійної пам'яті ЕОМ	До 500 Мб	До 300 Мб	До 1500 Мб	До 200 Мб	До 150 Мб
Використання оперативної пам'яті ЕОМ	До 32 гб	До 32 гб	До 32 гб	До 32 гб	До 16 гб
Зручний користувацький інтерфейс	+	-	+	+	+
Перелік можливих форматів вхідних даних	Власні формати, csv, xls	Власні формати, csv, xls	Власні формати, csv, xls	Власні формати, csv, xls	CUDA, csv

З таблиці можна зробити висновок що розроблена система прогнозування виграє у вартості, кількості оброблюваних джерел, можливості інтеграції у більші системи, використанні оперативної та постійної пам'яті обчислювальної машини, можновості розширення, але програє більшості аналогів у універсальності, гнучності, користувацькому інтерфейсі, переліку вхідних форматів даних, інтеграції із популярними програмними продуктами.

1.12 Вимоги до системи

Проаналізувавши наявні на ринку продукти і оглянувши стан проблемної області, сформуємо набір функціональних і нефункціональних вимог до системи.

1.12.1 Функціональні вимоги

До системи висуваються наступні функціональні вимоги:

- Можливість прогнозування часових рядів;
- Можливість аналізу тональності повідомлень соцмереж;

- Можливість аналізу тональності фінансових новин;
- Можливість ручного імпорту кількісних парламентів;
- Можливість збору кількісних параметрів із зовнішнього агрегатора;
- Можливість обробляти sitemap.xml сторінки сайтів з метою виділення

новин про певну компанію;

- Можливість зереження даних про тональність;
- Стійкість до проблеми зникаючого градієнту;
- Можливість авторизації і ліцензування;
- Можливість моніторингу фінансових новин;
- Можливість до сповіщення користувача про нові фінансові статті про певну

компанію у Slack;

1.12.2 Нефункціональні вимоги

До системи висуваються наступні не функціональні вимоги:

- Точність прогнозування тренду більше 57%;
- Час SQL-запиту до бази не більше 200 мілісекунд;
- Легка розширюваність;
- Дружній інтерфейс користувача;
- Кросплатформеність;
- Можливість до розширення списку сайтів-постачальників фінансових

новин;

На основі вмог до системи було побудовано діаграму варіантів використання. Діаграма варіантів використання наведена у додатку А.

Висновки до розділу

У розділі були розглянуті основні поняття у сфері прогнозування фінансових ринків, розглянуто специфіку прогнозування фінансових ринків, проведено огляд поточного стану проблемної області і огляд існуючих рішень на ринку.

Аналіз застосування традиційних методологій прогнозування показує, що вони не відповідають сучасним потребам ринку у прогнозуванні фінансових активів. Більшість методів розглядають вирішення задачі прогнозування лише з одного боку, беручи до уваги лише частину факторів, що впливають на зміну курсу фінансового активу і не враховують надважливі фактори, такі як реакцію людей у соціальних мережах і фінансові новини. У той час як фундаментальний аналіз розглядає зовнішні фактори, його застосування у системах прогнозування досить обмежене через складність реалізації і підтримки, або ж навпаки його реалізація досить неповна і орієнтується на зміну основних економічних і політичних індексів і звітів.

Існує потужний методологічний апарат для обробки часових рядів, методи авто регресії, ковзаючого середнього та інші. Проте економічні процеси керуються подіями, а часовий ряд (графік зміни курсу) зазвичай вже є наслідком цих подій.

Результатом дослідження є сформовані функціональні і не функціональні вимоги до системи, побудована діаграма варіантів використання. Для подільшого дослідженн і застосування у системі було обрано методи на основі часових рядів

Резюмуючи, можна сказати що існує гостра ринкова потреба у системі, яка буде суміщувати обробку часових рядів за допомогою нейромереж із аналізом додаткових факторів, притаманних для фундаментального аналізу. Для побудови такої системи необхідного розробити інструменти аналізу повідомлень у соцмережах із виявленням тональності повідомлення та фінансових новин, а також апарат аналізу економічних показників компанії. Ринкова потреба у такому інструменті підтверджує актуальність цієї магістерської дисертації.

2 МАТЕМАТИЧНІ МЕТОДИ ПРОГНОЗУВАННЯ ЧАСОВИХ РЯДІВ НА ФІНАНСОВИХ РИНКАХ

2.1 Визначення тензорних математичних об'єктів, ранг тензора

При прогнозуванні часових рядів, фінансових ринків та процесів використовуються математичні об'єкти, визначені в просторі одного, двох, трьох і більше вимірів. Ці математичні об'єкти мають відповідні аналоги у природі та техніці. Математично вони визначаються взаємозв'язаним набором чисел або функцій просторових координат[23].

Прикладами таких об'єктів є вектори, зокрема, вектори сили, швидкості, прискорення, напруженості електричного поля тощо. Вектори визначаються трьома окремими числами, або трьома окремими функціями просторових координат. Набір трьох чисел або функцій, що визначає вектор, змінюється при зміні системи координат, але сам математичний об'єкт — вектор — залишається незмінним. Він зберігає свою величину та орієнтацію у просторі (напрямок). Ці параметри допускають чітке фізичне трактування для більшості векторів.

При описі фізичних явищ використовуються математичні об'єкти, які представляють собою набір чисел або скалярних функцій у кількості більше трьох. Кількість функцій тісно пов'язана з розмірністю простору, в якому розглядається математичний об'єкт. Властивості цього об'єкта не залежать від вибраної системи координат, хоча при зміні системи координат значення чисел або скалярних функцій змінюється.

Такі об'єкти називаються тензорами, мають специфічну природу і визначаються в кожній точці простору. Тензори розглядаються у вибраній системі координат відносно конкретної точки простору з координатами $x_1, x_2, \dots, x_n$ і визначаються як деякі функції точки $Q(x_1, x_2, \dots, x_n)$. Ці функції являють собою пакет, складений із певної кількості скалярних функцій.

За визначенням математичні об'єкти у вигляді тензорів не залежать від вибраної системи координат $(x_1, x_2, \dots, x_n)$.

Розрізняють тензори нульового рангу (скаляри), першого рангу (вектори) та тензори більш високого рангу. Поряд з терміном ранг іноді вживають термін валентність.

Математичний об'єкт називається тензором (коваріантним) тоді, коли при зміні системи координат його компоненти перетворюються за законами перетворення одиничних векторів базисів старої та нової систем координат.

Кількісно тензори характеризуються функцією точки $Q(x_1, x_2, \dots, x_n)$, яка еквівалентна n^R скалярним функціям (R — ранг тензора). Наприклад, для тензора 2-го рангу ($R = 2$) в просторі трьох вимірів ($n = 3$) число скалярних функцій, які визначають функцію точки $Q(x_1, x_2, \dots, x_n)$, дорівнює $n^R = 3^2 = 9$.

Тензори 2-го і вищого рангів в залежності від характеру перетворення їх компонент можуть бути коваріантними, контраваріантними або змішаними.

Тензори першого рангу є векторами. Функція точки $Q(x_1, x_2, \dots, x_n)$, що відповідає тензору першого рангу, визначається пакетом із n (n — розмірність простору) скалярних функцій, які називаються компонентами вектора (тензора 1-го рангу). Для тривимірного простору маємо 3 функції.

Компоненти тензора першого рангу (вектора) змінюються при переході до нової системи координат. Схема, наведена на рисунку 2.1, ілюструє зміну компонент вектора при переході від однієї системи координат до іншої.

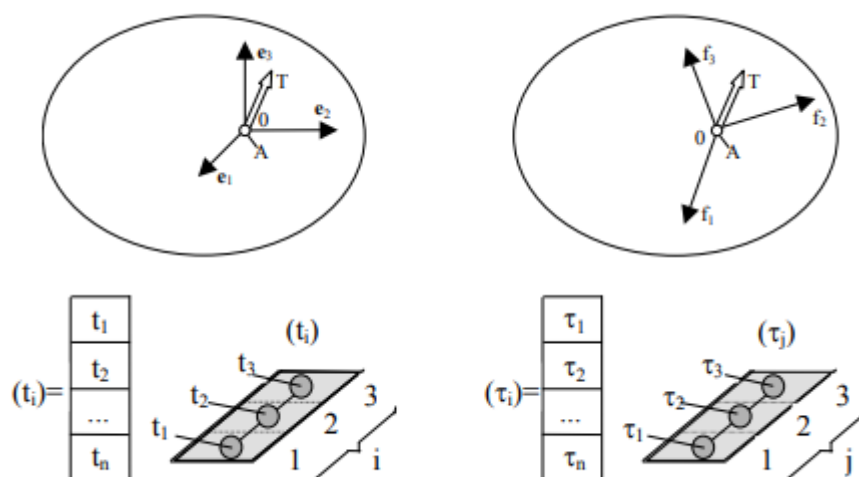


Рисунок 2.1 — Ілюстрація зміни компонент тензора першого рангу

Пакет скалярних функцій звичайно записують у вигляді таблиці, яка включає один стовпчик або один рядок. Таким чином, тензор 1-го рангу (вектор) можна записати у вигляді стовпчика чисел, прямокутної або просторової таблиці.

Вектор T визначений своїми трьома компонентами у першій системі координат $T = (t_i)$, а у другій системі координат — $T = ({}_j)$.

Вектор T як цілісний об'єкт не змінюється при зміні системи координат, але його компоненти змінюються за законом:

$$f_j^i Q(e_1, e_2, \dots, e_n) + f_j^i Q(x_1, x_2, \dots, x_n) = {}^*Q(x_1, x_2, \dots, x_n) {}^*T \quad (2.1)$$

де f_j^i — матриця переходу від системи координат з базисом e_1, e_2, e_3 до системи координат з базисом f_1, f_2, f_3 .

При зміні системи координат компоненти вектора змінюються, але деяка їх комбінація залишається незмінною (інваріантною). Ця величина, яка не залежить від зміни системи координат, називається інваріантом[23].

Як вказано раніше, вектори в залежності від перетворення їх компонент можуть бути двох типів: коваріантними або контраваріантними. Для коваріантних векторів використовуються позначення з індексом знизу, тобто t_i , а для контраваріантних — позначення з індексом зверху, тобто t^i .

2.2 Тензори другого рангу та їх графічна інтерпретація

Тензор 2-го рангу визначається як n^2 функцій точки і може бути представлений як:

- двічі коваріантний тензор, загальна компонента якого записується у вигляді t_{ij} ;
- двічі контраваріантний тензор, загальна компонента якого записується у вигляді t^{ij} ;

– змішаний тензор один раз коваріантний та один раз котраваріантний, загальна компонента якого записується у вигляді t_{ij} .

Тензори 2-го рангу визначені в просторі різних вимірів, починаючи з двовимірного простору ($n = 2, 3, 4, \dots$). При цьому індекси компонент тензора змінюються в межах $i, j = 1, 2, \dots, n$.

2.3 Декомпозиція Такера

Декомпозиція Такера — мультилінійне узагальнення добре відомої сингулярної декомпозиції матриць, що використовується у латентному семантичному аналізі [23]. У декомпозиції Такера тензор розкладається на ядерний тензор, який множиться на матриці по кожному з вимірів. Для тривимірного тензору X

$$X = I_1 \times I_2 \times I_3 = (I_1 \times R_1)(I_1 \times I_2 \times I_3) = G * Q \quad (2.2)$$

При цьому $P, Q, R \ll I, J, L$. Ядерний тензор G представляє стиснутий латентний варіант. На рисунку 2.2 зображено графічне представлення декомпозиції Такера.

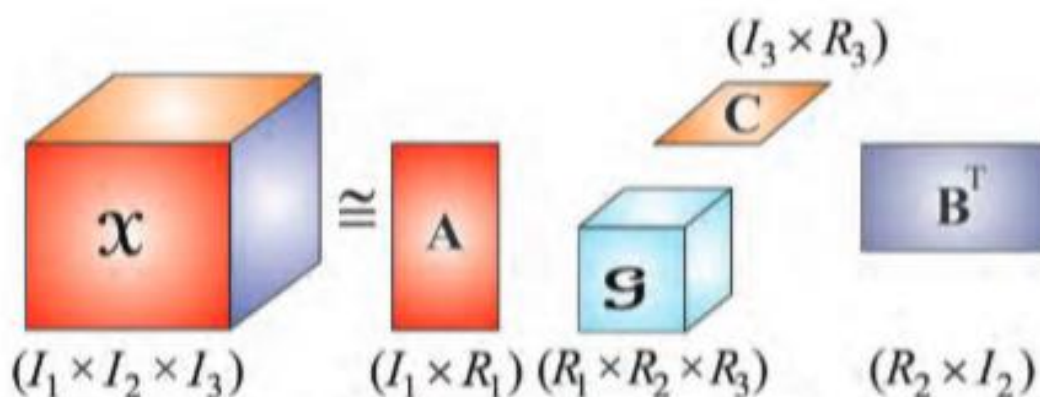


Рисунок 2.2 — Графічне представлення декомпозиції Такера

Декомпозиція Такера потрібна нам для того, щоб скоротити розмірність тензора, що призведе до підвищення точності прогнозування

2.4 Аналіз тональності повідомлень та новин про фінансові активи

Кількість інформації яка генерується за день у галузі фінансових ринків є завеликою для обчислення і аналізу у реальному часі. Кожної години з'являються десятки мільйонів нових твіттер-повідомлень, у два рази більше фейсбук-повідомлень і сотні тисяч статей регулярних видань різними мовами і близьмо пів-мільярда приватних повідомлень. З кожним роком темпи збільшення нової інформації у соціальних мережах зростають і, з великою ймовірністю, будуть зростати і надалі.

Якщо порівнювати соціальні мережи з точки зору доцільності застосування інформації взятої з них у нашій системі, можна зробити висновок, що Твіттер є найбільш прийнятним, незважаючи на те, що генерує чисельно менше інформації ніж Фейсбук, Інстаграм, іншими. Це обумовлено тим, що більшість користувачів Твіттер майже миттєво та у стислій формі реагує на нову для них інформацію. Це з одного боку полегшує задачу для системи аналізу, з іншого надає отриманим даним більшої актуальності. Вбудовані метадані, зручний для розробника зовнішній інтерфейс та відкритість роблять Твіттер майже ідеальним джерелом даних для нашого дослідження. Також перевагою Твіттер є те, що керівництво компанії з 2018 року проводить досить жорстку політику з фільтрації та блокування користувачів. Масово блокуються аккаунти неактивних користувачів, користувачів-ботів та користувачів що порушують користувацьку угоду. Це призвело до того, що станом на середину 2018 року інформацію у Тіттері генерують загалом живі активні користувачі, з чого випливає що ця інформація більш цінна для аналізу. Немає сенсу аналізувати повідомлення ботів та генераторів тональних повідомлень, оскільки така інформація є шкідливою для системи аналізу. Зауважимо, що сама система аналізу не має змоги виокремити повідомлення ботів від повідомлень живих людей.

Порівняння соціальних мереж і вибір одної потрібен через велику складність обробки даних з усіх можливих соціальних мереж через надвеликий об'єм даних що треба класифікувати, забрати та проаналізувати. При наявності великих обчислювальних потужностей та над швидкого з'єднання з мережею Інтернет доцільно обробляти всі можливі джерела. Порівняння соціальних мереж за критеріями доцільності застосування у нашому дослідженні наведено у Таблиці 2.1

Таблиця 2.1 — Порівняння соціальних мереж як джерел даних для прогнозування

	Соціальні мережі				
	Facebook	Youtube	Instagram	Twitter	LinkedIn
Загальна кількість користувачів	1.2 млрд	2.3 млрд	0.8 млрд	0.4 млрд	0.3 млрд
Щоденна кількість активних користувачів	215 млн	400 млн	112 млн	90 млн	57 млн
Кількість публікацій на день	4 млн	557 тис.	12 млн	8 млн.	1.1 млн
Наявність відкритого API повідомлень	+	-	+	+	+
Наявність метаданих повідомлень	-	-	-	+	-

Продовження таблиці 2.2

Кількість повідомлень від ботів (оціночно)	10 %	10%	10%	10%	5%
Вбудовані засоби аналізу тональності	-	-	-	-	-
Особливості контенту, який генерують користувачі соцмережі	Зазвичай манша кількість більших публікацій за об'ємом	Відео-контент	Слабке економічне ком'юніті, мало економічно го контенту.	Швидка реакція людей на події. Велике економічне ком'юніті	Сомережа для професіональної діяльності.

2.5 Методи дослідження тональності повідомлень

У дослідженні[24] проведено дослідження впливу тональності повідомлень користувачів Твіттер на зміну курсів акцій. Було доведено, що настрої користувачів Твіттер впливають на декілька біржових днів вперед, що корелює з часом, коли твіт активно читають інші користувачі. Варто зауважити, що вплив найбільший у той час, коли повідомлення було опубліковане і поступово згасає з часом. Точність кореляції настроїв користувачів і курсів активу залежить від зашумленості даних та якості словника (у випадку перевірки тональності за словником).

Виявлення настрою користувачів соціальних мереж є частковим випадком задачі обробки мови. Теоретично, побудова природно-мовного інтерфейсу для комп'ютерів — дуже приваблива мета. Перші системи, такі як SHRDLU, працюючи з обмеженими граматиками і використовуючи обмежений словниковий запас,

виглядали надзвичайно добре і перспективно, надихаючи цим своїх творців. Однак оптимізм швидко зник, коли ці системи зіткнулися зі складністю і неоднозначністю реального світу. Розуміння природної мови є AI-повною задачею, тому що розпізнавання живої мови вимагає величезних знань системи про навколишній світ і можливості з ним взаємодіяти. Саме визначення сенсу слова «розуміти» — одна з головних задач штучного інтелекту. А теж, для вирішення задачі розпізнавання тональності необхідно застосувати більш просту методологію, яку описано нижче.

Усі повідомлення поділимо на три класи: позитивні, нейтральні та негативні. Теоретично, повідомлення можна поділити на нескінченно велику кількість класів (наприклад, дуже позитивні, позитивні, позитивно-нейтральні, нейтральні, негативно-нейтральні, дуже негативні). Інтуїтивно здається, що це покращить якість виявлення тональності. Але у дослідженні [25] було доведено, що збільшення кількості класів тональності не призводить до збільшення кореляції тональності зі зміною стану предмету повідомлень, а навпаки, ускладнює обробку даних без покращення результаті аналізу або залишає результат незмінним.

Повідомлення будемо розглядати як кортеж з трьох змінних: метадані, тональність, автор повідомлення. Доцільно для деякого набору авторів, що впливають на думку читачів більше за інших, застосовувати мультиплікатор, який є кратним кількості читачів або визначати його емпірично. Незалежно від застосування мультиплікатору і методів виокремлення ознак, процедура аналізу складається з наступних етапів:

- збір даних: виконується за допомогою використання інтерфейсу постачальника даних. Попередня обробка набору даних;
- Вибір ознак;
- Класифікація;
- Визначення тональності;

2.6 Попередня обробка набору даних

Мета попередньої обробки — вилучення небажаних, незмістовних частин, вилучення небажаних інформаційних інверторів і притаманних людям сполучень, що можуть призвести до неправильного визначення тональності. Варто зауважити, що попередня обробка не може виявити сарказм (один з видів сатиричного викриття, уїдлива насмішка, вищий ступінь іронії) і іронію (сатиричний прийом, в якому справжній сенс приховує або суперечить явному сенсу.). Саркастичні та іронічні повідомлення зазвичай класифікуються помилково. Але все ж, це незначна частка від кількості сіх повідомлень.

Попередня обробка даних проходить у кілька етапів:

- Повідомлення приводиться до нижнього регістру (маленьких літер), видаляються небажані знаки пунктуації, допоміжні символи, керуючі символи;
- Визначається основа слова;
- Видаляється заперечення слова (назвичай частка 'ні');

Загалом всі ці етапи добре досліджені та легко реалізуються за допомогою готових бібліотек.

2.7 Виокремлення ознак

Мета цього етапу — представити дані після попередньої обробки у вигляді прийнятному для подальшої обробки. Існує два основних метода, які широко застосовуються у подібних задачах класифікації: виокремлення набору слів, метод N-грам[29].

2.8 Модель N-грамм

Загалом, N-грамма являє собою послідовність символі. Це може бути послідовність звуків, слогів, слів або букв. Зазвичай найчастіше зустрічається N-грамма як послідовність слів. Послідовність з двох послідовних елементів часто називають біграммою, послідовність з трьох елементів називається триграмма. Чотири і більше елементів зазвичай називають N-граммою, N замінюється на кількість послідовних елементів. В галузі обробки природної мови N-грамма

використовується в основному для прогнозування на основі статистичних моделей. N-граммна модель розраховує вірогідність останнього слова N-грамми, якщо відомі всі попередні. При використанні цього підходу для моделювання мови передбачається, що поява кожного слова залежить тільки від попередніх слів.

Другим застосуванням N-грамм є виявлення плагіату. Якщо розділити текст на кілька невеликих фрагментів, представлених N грамами, їх легко порівняти із іншим таким чином отримати ступінь збігу порівнюваних документів. N-грамми часто успішно використовуються для категоризації тексту та мови. Крім того, їх можна використовувати для створення функцій, які дозволяють отримувати знання з текстових даних.

2.9 Модель “мішок слів”

Модель тексту "мішок слів" являє собою сумативний набір (не системний), слів. Основний об'єкт моделі мішка слів — це слово, забезпечене єдиним атрибутом, частотою входження цього слова в масив[30].

В моделі тексту "мішок слів" враховується лише кількість введених конкретних слів в початковий текст, при цьому ігноруються: як порядок слів в повідомленні, так і морфологічні форми представлення слів. Текст довільної довжини перетворюється у вектор. При цьому формується словник всіх рецензій. Кожен елемент вектору, утвореного із тексту, це кількість входжень слова у початковий текст. Наведемо приклад. Нехай є словник {тварини, коти, собаки, також, миші, це, людина, павук, і}. Тоді речення “Коти це тварини, собаки це також тварини” буде мати вигляд {2, 1, 1, 1, 0, 2, 0, 0}. Лексика може бути досить об’ємною і може розростися до мільйонів символів. Тому доцільно обмежити її певним значенням (наприклад, 10000). Тоді кожна ознака буде мати розмірність 10000x1.

Модель тексту "мішок слів" була запропонована в 1975 році Дж. Солтоном, і в даний час є однією з найпоширеніших у різних областях лінгвістичних досліджень,

як правило, в якості основи для більш складних, перш за все "векторних моделей тесту".

Виокремлюють два основних типи атрибутів кортежу мішка слів:

— Частотні — кожне значення у векторі відповідає кількості входжень ознак в документ d ; тоді $(d) \in (0; +\infty)$;

— Бинарні (наявність / відсутність), кожне значення у векторі d бінарне (true / false або 0/1) і відображає факт наявності ознаки f_i в документі d . Тоді $(d) = \{0,1\}$;

У роботі розглянемо дві моделі N-грам — слів та символів. При порівнянні цих моделей виявляється, що саме N-грами символів мають більше переваг для вирішення задачі. Переваги N-грам символів:

— Можливість виправлення або ігнорування помилок. Відомо, що більшість користувачів пишуть повідомлення у соціальних мережах із орфографічними помилками. N-грами слів будуть пошкоджені від орфографічних помилок, у той час як N-грами символів — ні;

— Більш гнучкі можливості до збереження і формалізації;

Створення вектору ознак неможливе без створення вектору ваг ознак. Вага ознаки зазвичай відповідає кількості входжень, але може бути задана емпірично.

2.10 Методи, що базуються на використанні словників.

Тональність повідомлення залежить від тональності його окремих складових. Для аналізу складових повідомлення доцільно застосувати словники — списки лексем із вже відомою тональністю, яка дозволить віднести їх до певного класу. До переваг такого методу можна віднести високу надійність, до недоліків — необхідність формувати словники. Тому доцільно застосовувати вже створені словники, а не розробляти їх знову. Великі можливості для роботи зі словниками надає програмний пакет LIWC.

Важливо дослідити не тільки бінарну полярність лексеми, але силу цієї тональності, виражену численно. Для цього існує клас валентних словників. Для англійської мови була розроблена відкрита база даних тональних прикметників на

основі зв'язків понять з мережі WordNet. Основний алгоритм реалізації описано в роботі [30]. Для кожного значення емоційно-забарвленого прикметника приписується оцінка з нуля до одиниці по трьом шкалам:

- Pos (позитивність);
- Neg (негативність);
- Obj (об'єктивність);

Таким чином, сума всіх трьох оцінок рівна одиниці. Отримана оцінка візуалізована у вигляді трикутника з різнокольоровими вершинами. Насправді оцінка приписується цілому синсету, в якому з'являються одразу кілька прикметників за схожим значенням. Підхід класифікації синсетів дозволяє ефективно вирішити проблему збереження сенсу. Щоб використовувати такий словник на реальних даних, значення прикметників треба задавати самостійно. Тим не менш, далеко не всі прикметники мають великий розкид тональної оцінки залежно від сенсу, тому такий словник може виявитися ефективним інструментом при вирішенні задач аналізу тональності. Тому SentiWordNet активно застосовується в роботі з англійськими текстами.

Техніка створення англійського ресурсу, описана в роботі [31], ґрунтується на *semi-supervised learning* — фактично кортежі зі зв'язків синсетів в WordNet. Спочатку береться невелика кількість прикметників з явно вираженою позитивною або негативною тональністю. Потім ці множини поповнюються за рахунок додавання синонімічних синсетів тієї ж категорії або антонімічних синсетів з протилежного категорії. Решта синсетів потрапляють в клас Obj. На отриманій вибірці методами опорних векторів і Наївного Байєсовського класифікатору навчаються три класифікатори, які визначають позитивність, негативність і об'єктивність відповідно. Потім кожному класифікатору з найкращими параметрами пропонується визначити, чи не потрапили в навчальну вибірку синсети. Оцінка синсету виставляється як зважене оцінок всіх трьох класифікаторів.

2.11 Рекурентні нейронні мережі.

Я відомо, нейронні мережі прямого поширення використовують єдиний вхід, що надалі впливає на активації нейронів у інших шарах мережі. Через це виникає проблема — мережу неможливо навчити передбачати події, наступні елементи послідовностей, закінчення слів тощо. Для таких передбачень мережі не вистачає інформації про її попередні стани. Цю проблему легко вирішують рекурентні нейронні мережі. Циклічні рефлексивні внутрішні зв'язки рекурентної мережі дозволяють зберігати інформацію про попередні стани. При тому стан нейронів визначається не лише активацією у прихованих шарах, але й отриманими раніше активаціями того ж самого нейрону.

Рекурентна мережа відрізняється від нейронної мережі прямого поширення тим, що вона може бути розглянута як абір копій однієї нейронної мережі. При цьому кожна така копія пов'язана із наступною копією. Сама мережа RNN має ланцюгову структуру і часто застосовується до даних, що також мають ланцюгову або послідовну структуру.

Зауважимо, що об'єм оперативної пам'яті необхідний для запуску мережі пропорційний кількості елементів у послідовності що оброблюється. Послідовність, що складається з N елементів буде представлена у пам'яті як N векторів. Проте розміри матриці ваг W не змінюється пропорційно кількості елементів у послідовності.[32] Так для рекурентної мережі що складається з K рекурентних шарів розмір матриці завжди буде дорівнювати $K \times K$ незалежно від кількості елементів у оброблюваній послідовності.

2.11.1 Різниця між рекурсивними та рекурентними мережами

Самоповторення рекурентних нейронних мереж відбувається у часі. Нехай необхідно спрогнозувати наступний елемент послідовності ("П", "Р", "И", "В". "І"). Всю послідовність від першого до останнього символу оброблює один і той самий нейрон, що має рефлексивний зв'язок із самим собою.

Так, на першому етапі буде подано символ П, на другому Р. На третьому Нейрон буде зберігати свій стан на минулому кроці і, зрештою, має правильно

вказати на наступний символ — “Т”. Процес прогнозування наступного символу зображено на рисунку 2.3.

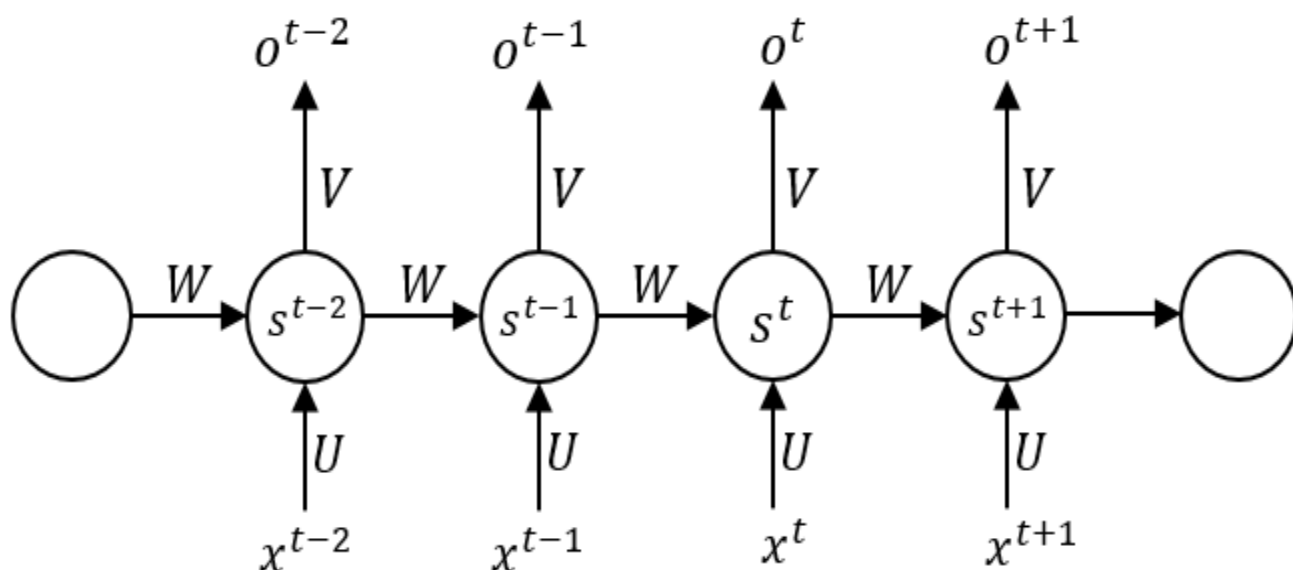


Рисунок 2.3 — Схема нейрону, що оброблює послідовність і прогнозує наступний символ послідовності

Застосовуючи подібний підхід важливо розуміти різницю між рекурентною та рекурсивною мережею. Взагалі рекурентна нейрона мережа є частковим випадком рекурсивної. Рекурентна мережа має спільні ваги для однієї розмірності на всій довжині оброблюваної послідовності. Рекурсивна мережа так само має спільні ваги для кожного вузла. З цього випливає, що всі ваги будуть однаковими для певного нейрону. Головна відмінність у двох типах мереж полягає у галузях, де доцільне їх застосування. Рекурсивні мережі застосовуються для генерації послідовностей. Рекурентні — для обробки дерев і прогнозування часових рядів.

2.11.2 Обмеження рекурентних нейронних мереж

Однією з основних проблем, яка виникає при застосуванні класичних рекурентних нейронних мереж є проблема зникаючого градієнту. Розглянемо випадок, коли може виникнути така проблема.

Нехай наша мережа навчена для виконання задачі прогнозування мовної послідовності — прогнозує наступне слово використовуючи всі попередні. Інтуїтивно здається, що така задача добре вписується в концепцію застосування класичних рекурентних мереж і буде вирішена легко. Так і є у тому випадку, коли між елементами послідовності у яких є контекстний зв'язок коротка відносна відстань у послідовності. Наприклад, послідовність “Коти споживають рибу, собаки споживають м'ясо” містить короткі контекстні зв'язки, найбільша довжина контекстного зв'язку — 2, наприклад між “Коти” і “рибу”. Такі послідовності передбачити зазвичай досить легко.

Але практика показує що більшість повідомлень містить достатньо довгі (3 і більше) відстані між словами що мають контекстний зв'язок. Наприклад, розглянемо повідомлення “Коти — мої улюблені тварини, тому коли я зможу я обов'язково куплю собі кота”. У цьому прикладі є контекстний зв'язок довжиною у 12 слів. Проміжок між словами що пов'язані значно ширший ніж у минулому прикладі.

Виявилось, що зі збільшенням довжини контекстного зв'язку рекурентну мережу неможливо начити зв'язувати дані. Цей приклад ілюструє проблему навчання через нестійкий градієнтний спуск. Попередні шари навчаються порівняно повільно і градієнт при зворотньому поширенні поступово зменшується. З цього випливає, що якщо мережа навчається довго, стійкість градієнту буде зменшуватися

$$\left\| \frac{\partial h_t}{\partial h_k} \right\| = \left\| \prod_{j=k+1}^t \frac{\partial h_j}{\partial h_{j-1}} \right\| \leq (\beta_W \beta_h)^{t-k} \quad (2.3)$$

У момент коли значення градієнту стає завеликим, воно стає причиною переповнення, яке досить легко виявити під час навчання нейронної мережі. Це переповнення пов'язано із проблемою градієнтного вибуху. Протилежна проблема — проблема зникаючого градієнта. Вона виникає коли градієнт стає риним нулю та зазвичай значно важче виявляється під час навчання. При цьому якість навчання

значного погіршується і чим далі шар, тим більше це впливає на якість його навчання.

Теоритично рекурентні нейромережі можуть обробляти довгі контекстні залежності. Для цього потрібно корегувати параметри вручну, або розробити окремий модуль що взаємодіє із нейронною мережею та регулюю її. Але найпростіше і найкраще вирішення проблеми — додавання модулю довгої короткострокової пам'яті до нейронної мережі.

2.12 Довга короткострокова пам'ять (LSTM)

Нейронні мережі із використанням довгої короткострокової пам'яті здатний запам'ятовувати контекстні зв'язки як для елементів із довгою відстанню між ними, так і для елементів що розташовані поряд один з одним. Зауважимо, що сам LSTM-модуль не використовує функцію активації всередині своїх компонентів. Через це значення комірки не зникає із часом, а градієнт не може зникати під час навчання нейронної мережі методом зворотнього поширення помилки.

Метою розробки модулів довгої короткострокової пам'яті було формування більш стійкої пам'яті мережі, що значно би спростило фіксацію довгих контекстних залежностей. Загальним для всіх рекурентних мереж є те, що вони складені з ланцюга повторюваних блоків. Цей блок зазвичай має дуже просту структуру. Зазвичай, це єдина $\tanh$ -функція у стандартних рекурентних мережах.

Комірка довгої короткострокової пам'яті також має таку саму ланцюгову структуру, але комірка має вже не один, а чотири блоки що впливають один на одного. Цю комірку зображено на рисунку 2.5

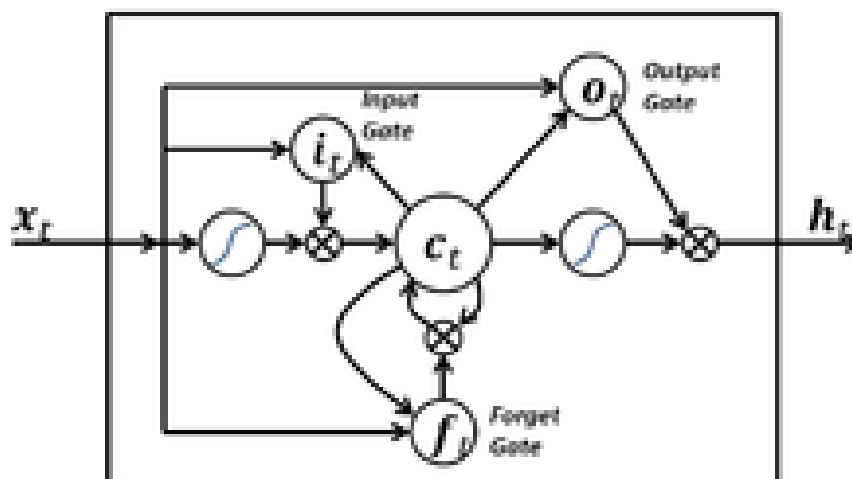


Рисунок 2.5 — структура комірки довгої короткострокової пам’яті

LSTM-модуль має дуже корисну здатність до обробки послідовностей із довгостроковими залежностями: здатність додавати інформацію у стан комірки. При цьому цей процес регулюється “фільтрами” — регуляторами, які відповідають за процес забування чи зберігання інформації у комірці.

Комірка пам’яті LSTM зберігає значення для довгих і коротких періодів часу. Це досягається за рахунок активаційної функції для кожної комірки пам’яті. Однією із переваг LSTM є те, що вона позбавлена ефекту зникнення градієнту. Схема нейрону LTSM наведена на рисунку 2.6

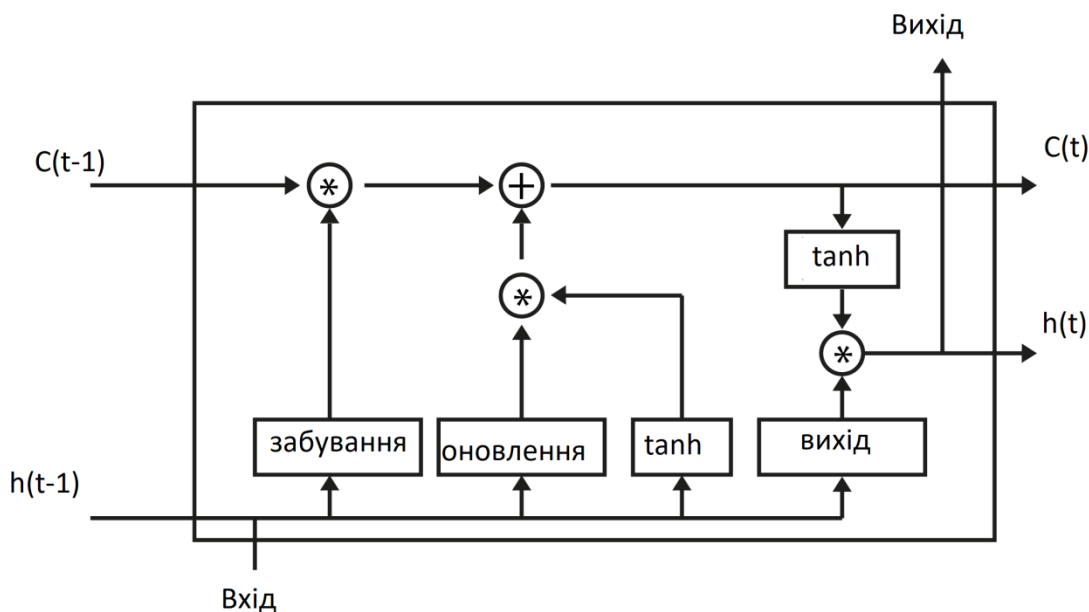


Рисунок 2.6 — Комірка LSTM

Головний компонент LSTM — це стан ячейки (cell state) — горизонтальна лінія, що проходить по верхніх частинах схеми. Інформація у ній може зберігатися без змін. Проте, LSTM може видаляти інформацію зі стану ячейки; цей процес регулюється фільтрами (gates). Фільтри дозволяють пропускати інформацію на підставі деяких умов. Вони складаються з шару сигмоїдальної нейронної мережі і операції множення.

По-перше треба визначити, яку інформацію можна викинути зі стану ячейки. Це рішення приймає сигмоїдальний шар, "шар фільтра забування" (forget gate layer). Наступний крок — вирішити, яка нова інформація буде зберігатися в стані ячейки. Цей етап складається з двох частин. Спочатку сигмоїдальний шар що зветься "шар вхідного фільтра" (input layer gate) визначає, які значення слід оновити. Потім $\tanh$ -шар будує вектор нових значень-кандидатів. Нарешті, потрібно вирішити, яку інформацію ми хочемо отримувати на виході. Вихідні дані будуть засновані на нашому стані ячейки, до них будуть застосовані деякі фільтри. Спочатку ми застосовуємо сигмоїдальний шар, який вирішує, яку інформацію зі стану ячейки ми будемо виводити. Потім значення стану осередку проходять через $\tanh$ -шар, щоб отримати на виході значення з діапазону $[-1; 1]$, і перемножуються з вихідними значеннями сигмоїдального шару, що дозволяє виводити тільки необхідну інформацію.

2.12.1 Архітектура комірки довгої короткострокової пам'яті.

Для генерації нової пам'яті використовується вхідний сигнал x_t та минулий прихований стан h_{t-1} . Разом вони утворюють нову пам'ять c_t , яка ураховує як минули стан, так і нові дані із деякими коефіцієнтами, за які відповідають три фільтри-регулятори.

Фільтр входу

Фільтр входу перевіряє чи варто зберігати нове повідомлення, перед тим як створити новий стан комірки. Він використовує як новий сигнал так і попередній

прихований стан. Результатом його роботи є коефіцієнт i_t який впливатиме та майбутній стан комірки.

Фільтр забування

Мета фільтра забування — оцінити кількісно наскільки корисний минулий блок пам'яті для обчислення минулого блока пам'яті. Результатом його роботи є коефіцієнт f_t який також впливатиме та майбутній стан комірки.

Блок пам'яті

Блок пам'яті відповідає за стан комірки. Від виконує дві обчислювальні операції. Спочатку він забуває відповідно до коефіцієнта f_t минулий стан ячейки c_{t-1} . Це потрібно для того, щоб оцінити наскільки нова інформація важлива для комірки пам'яті і чи треба її перезаписувати і замінювати. Потім така сама операція застосовується до показника i_t і блоку нової пам'яті c_t .

Потім результати обчислень першого і другого етапі сумуються і в результаті отримуємо новий блок пам'яті c_t .

Фільтр виходу

Фільтр виходу відповідає за розділення блоку пам'яті та прихованого стану комірки. Як було сказано вище, блок пам'яті c_t містить велику кількість інформації, але не всю цю інформацію потрібно зберігати у прихованому стані. Прихований стан використовується у кожному фільтрі довгої короткострокової пам'яті, тому фільтр виходу оцінює, які саме дані з пам'яті c_t мають бути примутніми у прихованому стані h_t

Математично, робота складових комірки виглядає так:

Фільтр входу:

$$i_t = \sigma(W^{(i)}x_t + U^{(i)}h_{t-1}) \quad (2.4)$$

Фільтр забування:

$$f = \sigma(W^{(f)}x_t + U^{(f)}h_{t-1}) \quad (2.5)$$

Фільтр виходу:

$$f = \sigma(W^{(o)}x_t + U^{(o)}h_{t-1}) \quad (2.6)$$

Новий блок пам'яті

$$c_t = \tanh(W^{(c)}x_t + U^{(o)}h_{t-1}) \quad (2.7)$$

Блок пам'яті і вихід:

$$c_t = f_f * c_{t-1} + i_t * c_1) \quad (2.8)$$

$$h_t = o_t * \tanh(c_t) \quad (2.9)$$

2.12.3 LSTM-архітектури для задачі прогнозування часових рядів

У дослідженні компанії Google [33] розглядалася можливість покращення архітектури довгої короткострокової пам'яті. Було використано тисячі надпотужних ЕОМ і перевірено сотні варіацій архітектури. Результатом дослідження є те, що класична структура нейронної мережі довгої короткострокової пам'яті загалом краще справляється із задачею прогнозування ніж її варіації за критерієм точності прогнозу. Однак, було виявлено варіацію архітектури яка прогнозує деякі нестационарні часові ряди краще за класично. Це архітектура із рекурентним вузлом.

Особливістю архітектури із вентиляним рекурентним вузлом є те, що фільтри мають доступ до стану комірки. Математично, це виглядає так:

Фільтр входу:

$$i_t = \sigma(W^{(i)}x_t + U^{(i)}h_{t-1} + b_i) \quad (2.10)$$

Фільтр забування:

$$f = \sigma(W^{(f)}x_t + U^{(f)}h_{t-1} + b_f) \quad (2.11)$$

Фільтр виходу:

$$f = \sigma(W^{(o)}x_t + U^{(o)}h_{t-1} + b_o) \quad (2.12)$$

Новий блок пам'яті

$$c_t = \tanh(W^{(c)}x_t + b_c) \quad (2.13)$$

Блок пам'яті і вихід:

$$c_t = f_f * c_{t-1} + i_t * c_1 \quad (2.14)$$

$$h_t = o_t * c_t \quad (2.15)$$

Рекурентний вузол поєднує вхідний фільтр та фільтр забування у новий тип фільтру — фільтр оновлення, а також поєднує стан комірки з прихованим станом. Комірку із рекурентним вентилям оновлення зображено на рисунку 2.7

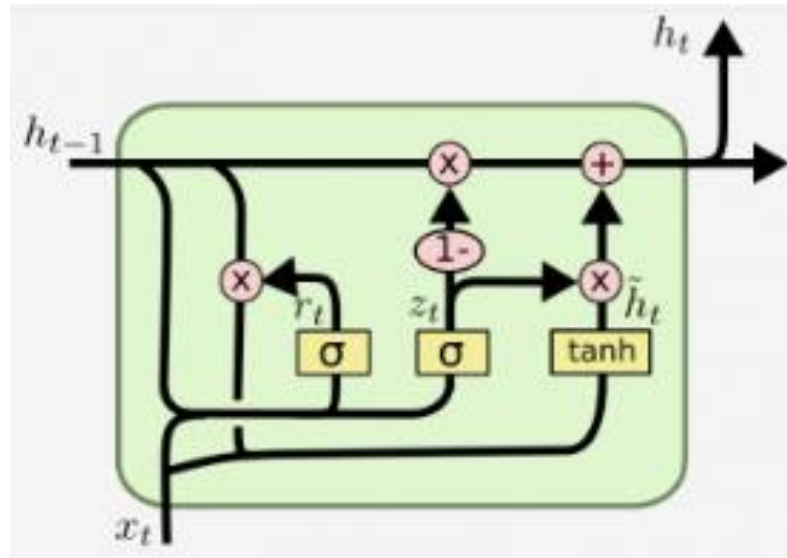


Рисунок 2.7 — комірка із вентилям рекурентним вузлом.

Математично робота вентиля оновлення виглядає так:

Фільтр оновлення:

$$z = \sigma(W^{(z)}x_t + U^{(z)}h_{t-1}) \quad (2.16)$$

Фільтр видалення:

$$r = \sigma(W^{(r)}x_t + U^{(r)}h_{t-1}) \quad (2.17)$$

Новий блок пам'яті

$$h = \tanh(W^0 x_t + U^0 h_{t-1} * r) \quad (2.18)$$

Прихований стан:

$$h_t = (1 - z_t) * h_t + z_t * h_{t-1} \quad (2.19)$$

2.13 Алгоритм скорочення розмірності тензору.

В основі алгоритму скорочення розмірності тензору лежить розкладання Такера, яке було опиано вище, та метод головних компонент. Метод головних компонент є динм із основних алгоритмів скорочення розмірності із втратою мінімальної кількоти даних.

В результаті розкладання такера, на кожній ітерації будемо отримувати 1 головний тензор і 3 матриці. Метод головних компонент працює в першу, чергу з матрицями.

Виконавши циклічне розкладання тензору за методом такера, отримаємо $N \times 3$ матриць, де N — початкова розмірність тензору. Застосуємо метод головних компонент для опорних матриць і ітеративно відновимо тензор.

Роботу алгоритму зображено у додатку А

Висновки до розділу

У даному розділі були розглянуті основні математичні підходи до складових системи прогнозування фінансових активів. Детально висвітлені алгоритми аналізу тональності тексту. Серед двох класів методів аналізу тональності повідомлень, а саме методів машинного навчання та методів заснованих на словниках, вибір було

зроблено на користь методів заснованих на словниках через простішу реалізацію обробки повідомлень соціальних мереж а фінансових новин.

Було розглянуто специфіку застосування нейромереж у задачі прогнозування фінансових активів. Через ланцюговий характер вхідних даних, вибір було зроблено на користь рекурентних нейромереж. Було розглянуто проблему зникнення контекстних зв'язків у рекурентних нейромережах і запропоновано механізм її вирішення — додавання LSTM-модулю до комірок рекурентних нейромереж. Для магістерської дисертації була обрана саме архітектура із довгою короткостроковою пам'яттю через те, що при прогнозуванні фінансових активів контекстні залежності у часових рядах зазвичай розташовані далеко одна від одної і застосування звичайних рекурентних нейромереж не враховувало би їх.

Було представлено алгоритм прогнозування фінансових активів, що використовує тензорне розкладання Такера із подальшою перебудовою тензору. Цей покликаний підвищити точність прогнозування за рахунок підвищення якості вхідних даних і зменшення їх надлишковості.

В результаті аналізу математичних методів, ндо застосування у системі були вибрані метод мішка слів, декомпозиція Такера та нейронні мережі із довгою короткостроковою пам'яттю. Результати аналізу показали, що нейронні мережі із довгою короткостроковою пам'яттю ільш придатні до задачі прогнозування фінансових ринків через здатність до роботи із довгими контекстними залежностями та відсутності проблеми зникнення градієнту.

3 РОЗРОБКА АРХІТЕКТУРИ СИСТЕМИ ПРОГНОЗУВАННЯ

Розроблений програмний продукт складається з п'яти незалежних модулів, а саме:

- модуль збору повідомлень із соціальних мереж та аналіз їх тональності;
- модуль збору фінансових новин та аналіз їх тональності;
- модуль обробки кількісних даних і показників роботи компанії-об'єкту прогнозування;
- модуль обробки тензору;
- неймережевий модуль;

Загальна архітектура системи зображена на рис 3.1

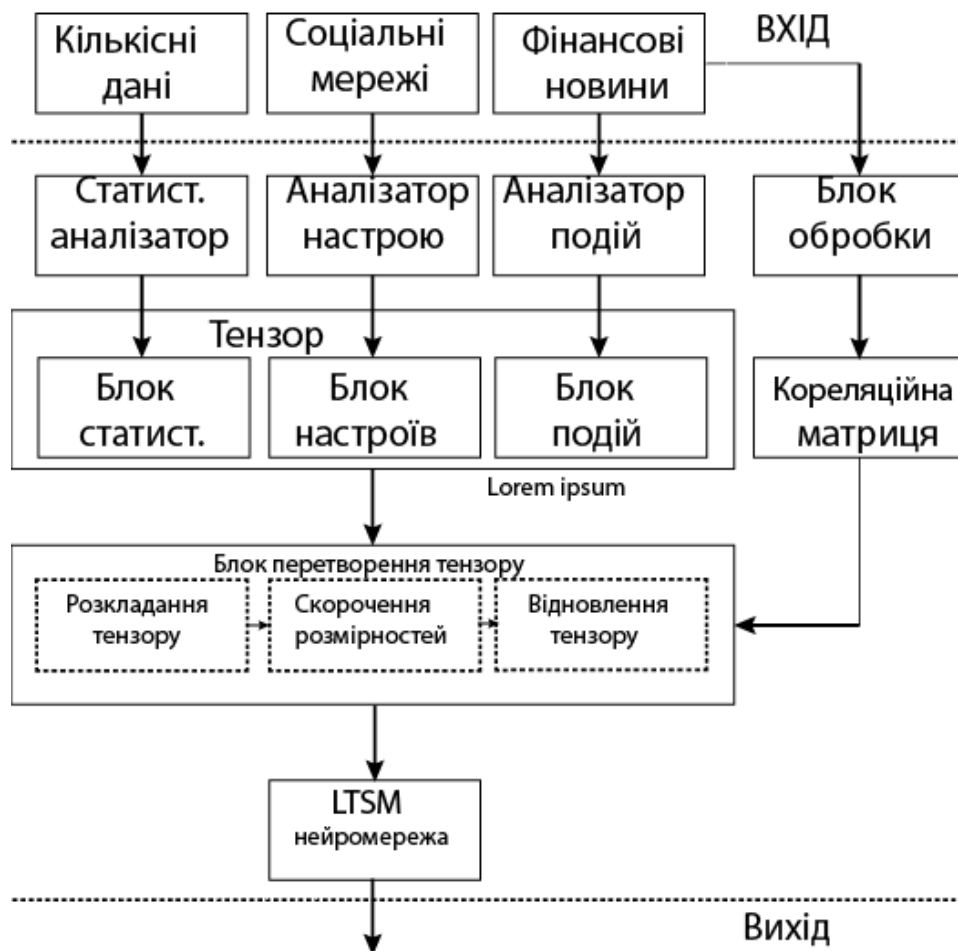


Рисунок 3.1 — Загальна архітектура системи

Така архітектура відповідає принципу єдиної відповідальності із набору принципів SOLID. Кожен модуль відповідає за окремий блок функціональності, має вхідні і вихідні дані узгоджені з іншими модулями. Розбиття системи на модулі

згідно принципу єдиної відповідальності значно спрощує тестування системи. Діаграма класів системи наведена у Додатку Б, діаграма пакетів системи наведена у додатку В.

3.1 Модуль для збору і аналізу даних з соціальних мереж

У модулі застосовується метод “мішок слів” та наївний класифікатор Басса, описані у розділі 2. Діаграма послідовностей модуля збору і аналізу даних з соціальних мереж наведена у Додатку Г

Модуль повинен задовольняти наступні вимоги:

- збір повідомлень із соціальної мережі, вибір якої був обґрунтований у розділі 2. При виборі необхідних повідомлень використовується механізм метаданих;
- Приведення повідомлення до нижнього регістру (маленьких літер), видалення небажаних знаків пунктуації, допоміжних символів, керуючих символів;
- Визначення основи слова;
- Видалення заперечення слова;
- Виокремлення ознак згідно методу мішку слів;
- Класифікація повідомлень;
- Визначення тональності;

3.1.1 Вибір мови програмування

Мову програмування було обрано за наступними критеріями:

- Сучасність і регулярна підтримка з боку розробників мови;
- Можливість розроблювати і запускати застосування у різних операційних системах (кросплатформеність);
- Наявність готових модулів для аналізу тональності тексту і бібліотек для вирішення задачі класифікації;

Згідно з перерахованими критеріями було обрано мову програмування Python

3.6. Ця мова не тільки стрімко розвивається і підтримується основними

операційними системами, а й має потужні вбудовані інструменти для аналізу тональності тексту і вирішення задач класифікації.

3.1.2 Бібліотки, що застосовуються у модулі.

Модуль використовує більше десяти вбудованих бібліотек та бібліотек сторонніх розробників, проте опишемо основні з них.

У процесі збору повідомлень, а саме у перетворенні html-коду на словники і структури мови для виокремлення повідомлень використовується бібліотека BeautifulSoup 4.

Beautiful Soup 4 — це парсер для синтаксичного розбору файлів HTML / XML, написаний на мові програмування Python, який може перетворити навіть неправильну розмітку в дерево синтаксичного розбору. Він підтримує прості і природні способи навігації, пошуку та модифікації дерева синтаксичного розбору. У більшості випадків роботи пов'язаною із обробкою веб-контенту він може допомогти заощадити години і дні роботи.

Для роботи із Twitter була використана бібліотека Python Social Auth, а класифікація була виконана за допомогою алгоритму vader, що реалізований у бібліотеці vaderSentiment.

Для збереження даних використовувалася система управління базами даних SQLite3. Вибір було зроблено через її можливості до легкого встановлення та налаштування та відсутності у потребі додаткових зовнішніх модулів.

3.1.3 Джерела повідомлень

Для отримання даних з мережі Twitter використовується Twitter Streaming API. Twitter Streaming API дозволяє безкоштовно отримувати потік фільтрованих записів з Twitter. Це дуже зручний механізм, який дозволяє отримувати повідомлення згідно з метаданими, що значно спрощує пошук потрібних для дослідження повідомлень.

3.1.4 Архітектура модуля.

Модуль складається з наступних класів.

- Клас `twitterScarper`. Виконує збір повідомлень за заданими метаданими;
- Клас `twitterParser`. Перетворює html-код на набір повідомлень. Виконує первісну обробку повідомлень;
- Клас `messagePreparer`. Виконує визначення основи слова, видалення заперечення слова та попередню обробку повідомлень до виду, придатного для подальшого аналізу;
- Клас `worldBagAnalyzer`. Виконує перетворення набору повідомлень згідно алгоритму “мішок слів”;
- Клас `Semantic Analyzer`. Визначає тональність повідомлень і формує дані, що будуть придатні до обробки у тензорі;

Діаграма класів модуля зображена на рисунку 3.2

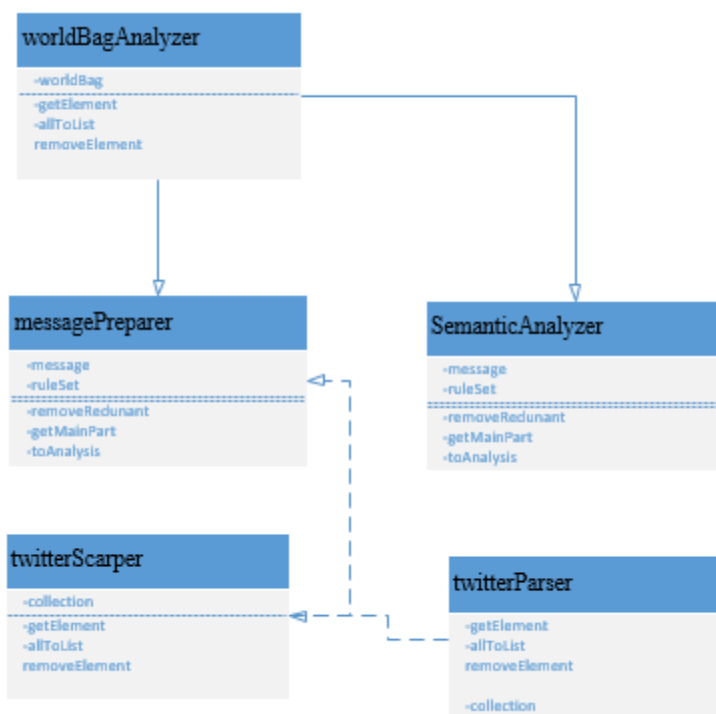


Рисунок 3.2 — Діаграма класів модулю обробки даних соціальних мереж

3.2 Модуль для збору і аналізу фінансових новин

Як було зазначено вище, обробка фінансових новин має декілька принципових відмінностей у процесі збору даних. Так при обробці повідомлень соціальних мереж ми працюємо із одним сайтом — Twitter. Або ще з декількома, якщо ми додамо, наприклад, Facebook або Вконтакте. У будь якому разі це робота із відомими API обмеженої кількості сайтів. У обробці фінансових новин так бути не може. Сайтів, що спеціалізуються на фінансових новинах тисячі, не всі з них мають відкритий та безкоштовний API. Не можна спиратися лише на новини декількох сайтів. Але кожен сайт зазвичай містить сторінку SITEMAP.XML, де перераховані всі наявні сторінки цього сайту. Корисною властивістю цього списку є те, що кожна сторінка містить також і дату публікації на SITEMAP.XML. Зазвичай для кожної новини створюється окрема сторінка, яка має інформативну назву у своїй адресі. Зазвичай у цій адресі міститься назва компанії, акції якої ми хочемо спрогнозувати. Наприклад, зрозуміло що сторінка “<https://ru.tsn.ua/groshi/akcii-apple-upali-posle-prezentacii-v-ssha-novyh-iphone-1216191.html>” містить новини про акції компанії Apple Inc. Доцільно вибрати такі сторінки, завантажити їх, та обробляти їх текст так само, як і у модулі обробки повідомлень соціальних мереж. У системі запропонований алгоритм збору фінансових новин, який містить наступні кроки:

- Створення списку сайтів які будемо аналізувати;
- Завантаження списку url-адрес сторінок кожного веб-сайту (сторінка sitemap.xml) зі списку формованого у пункті 1 і оновлення в режимі постійного моніторингу;
- Вибір тих URL-адрес, у складі яких є назва об’єкту прогнозування.
- Завантаження тексту сторінки;
- Обробка тексту так само як і у модулі обробки соцмереж.
- Визначення тональності тексту;

Діаграма послідовностей модулю наведена у Додатку Д

Цей модуль повинен задовольняти наступним вимогам:

- Можливість створення списку сайтів, фінансові новини яких будемо аналізувати;
- Можливість завантаження сторінки sitemap.xml цих сайтів;
- Парсинг sitemap.xml і виявлення тих нових, які стосуються об'єкту прогнозування;
- Завантаження коду цих обраних сторінок і виокремлення текстових повідомлень із них;
- Приведення повідомлення до нижнього регістру (маленьких літер), видалення небажаних знаків пунктуації, допоміжних символів, керуючих символів;
- Можливість визначення основи слова;
- Можливість видалення заперечення слова;
- Можливість виокремлення ознак згідно методу мішку слів;
- Можливість класифікації повідомлень;
- Можливість визначення тональності;
- Можливість сповіщати користувача про появу нових фінансових новин про об'єкт прогнозування;

3.2.1 Вибір мови програмування

При виборі мови програмування були розглянуті такі самі критерії як і для модулю аналізу повідомлень соціальних мереж:

- Сучасність і регулярна підтримка з боку розробників мови;
- Можливість розроблювати і запускати застосування у різних операційних системах (кросплатформеність);
- Наявність готових модулів для аналізу тональності тексту і бібліотек для класифікації;

Згідно з перерахованими критеріями було обрано мову програмування Python 3.6. Ця мова не тільки стрімко розвивається і підтримується основними операційними системами, а й має потужні вбудовані інструменти для аналізу тональності тексту і вирішення задач класифікації.

3.2.3 Бібліотки, що застосовуються у модулі.

Для створення графічного інтерфейсу використовувався пакет PyQt5. PyQt5 — це набір Python бібліотек для створення графічного інтерфейсу на базі платформи Qt5 від компанії Digia. Він доступний для Python 2.x і 3.x. Бібліотека Qt є однією з найпотужніших бібліотек GUI (графічного користувацького інтерфейсу)

Інші використані бібліотеки такі самі як і у модулі обробки повідомлень соціальних мереж.

Для отримання фінансових новин у нашій системі користувач має змогу власноруч визначити список сайтів-постачальників фінансових новин, або додати до списку, що був сформований під час дослідження.

Роботу модулю обробки новин у користувацькому інтерфейсі зображено на рисунку 3.4

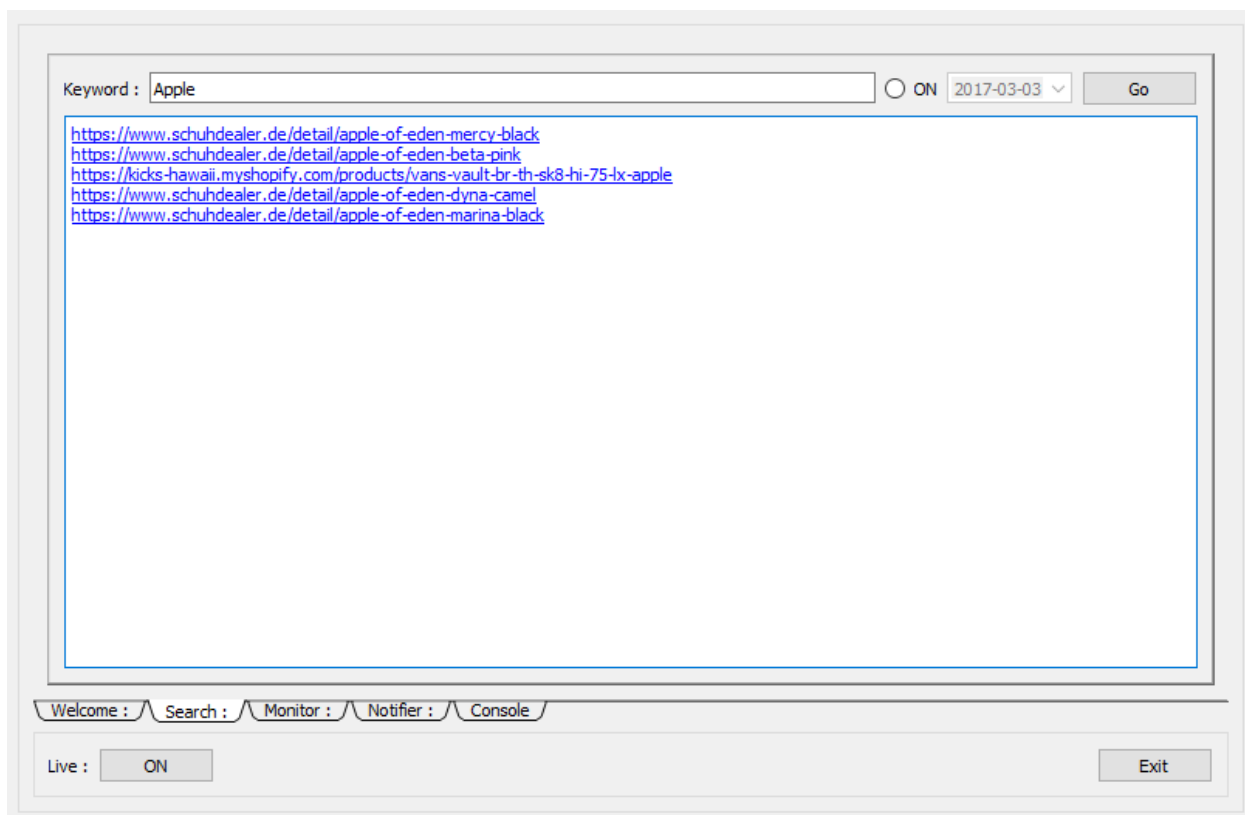


Рисунок 3.4 — Робота модулю обробки новин у користувацькому інтерфейсі

Список сайтів, які були включені до списку у цьому дослідженні:

- European Central Bank;
- Bank of England;
- FT.com;
- Bloomberg;
- Guardian;
- Bureau of Labor Statistics;
- Money Week;
- UK Statistics Authority;

3.2.4 Архітектура модуля.

Діаграма класів модуля зображена на рисунку 3.5

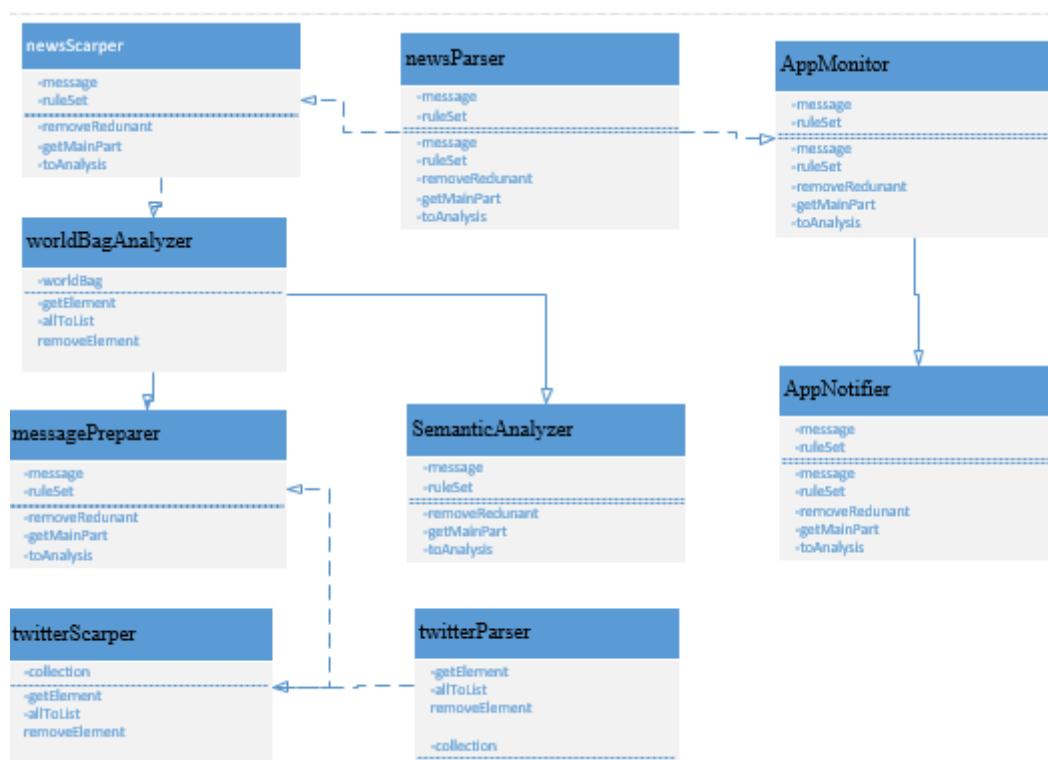


Рисунок 3.5 — Діаграма класів модулю обробки даних фінансових новин

Модуль складається з наступних класів:

- Клас newsScarper. Виконує збір новин за заданим критерієм пошуку;
- Клас appMonitor. Виконує моніторинг новин за заданим критерієм;

- Клас `appNotifier`. Повідомляє користувачу в Slack про нову фінансових новин про об'єкт прогнозування;
- Клас `newsParser`. Перетворює html-код на набір повідомлень. Виконує первісну обробку повідомлень;
- Клас `messagePreparer`. Виконує визначення основи слова;
- видалення заперечення слова та попередню обробку повідомлень до виду, придатного для подальшого аналізу;
- Клас `worldBagAnalyzer`. Виконує перетворення набору повідомлень згідно алгоритму “мішок слів”;
- Клас `Semantic Analyzer`. Визначає тональність повідомлень і формує дані, що будуть придатні до обробки у тензорі;

Наведемо приклад. Нехай є словник {тварини, коти, собаки, також, миші, це, людина, павук, і}. Тоді речення “Коти це тварини, собаки це також тварини” буде мати вигляд {2, 1, 1, 1, 0, 2, 0, 0}. Лексика може бути досить об'ємною і може розростися до мільйонів символів. Тому доцільно обмежити її певним значенням (наприклад, 10000). Тоді кожна ознака буде мати розмірність 10000x1.

4.3 Модуль для збору і аналізу кількісних даних

Цей модуль повинен задовольняти наступним вимогам:

- збір кількісних даних про об'єкт прогнозування;
- можливість імпорту даних про об'єкт прогнозування у форматі CSV;
- нормалізація і попередня обробка кількісних даних;
- підготовка даних до подальшої роботи із тензором;

Діаграма послідовностей модулю наведена у Додатку Е, алгоритм роботи модулю наведений у додатку Ж

4.3.1 Бібліотки, що застосовуються у модулі.

Для аналізу даних у модулі використовуються пакети `Numpy` та `Pandas` мови Python 3.6. Якщо дуже коротко, то `pandas` — це бібліотека, яка надає дуже зручні з

точки зору використання інструменти для зберігання даних і роботи з ними. Якщо ви займаєтеся аналізом даних або машинним навчанням і при цьому використовуєте мову Python, то ви просто зобов'язані знати і вміти працювати з pandas.

Для збереження даних використовувалася система управління базами даних SQLite3. Вибір було зроблено через її можливості до легкого встановлення та налаштування та відсутності у потребі додаткових зовнішніх модулів.

Дані можуть бути як завантажені із агрегатору економічних даних CUBA через відкритий та безкоштовний інтерфейс, так і завантажені вручну. Головне — використовувати один і той самий формат даних. Приклад формату даних наведений у Додатку В.

Модуль складається з наступних класів:

- Клас `dataMonitor`. Виконує збір кількісних даних з агрегатору економічних даних CUBA;
- Клас `dataImporter`. Виконує ручний імпорт даних із CSV-файлу у фаловій системі користувача;
- Клас `dataNormalizer`. Виконує нормалізацію на первісну підготовку отриманих кількісних даних;
- Клас `tensorBuiler`. Формує дані до обробки у тензорі;

Діаграма класів модулю зображена на рисунку 3.6

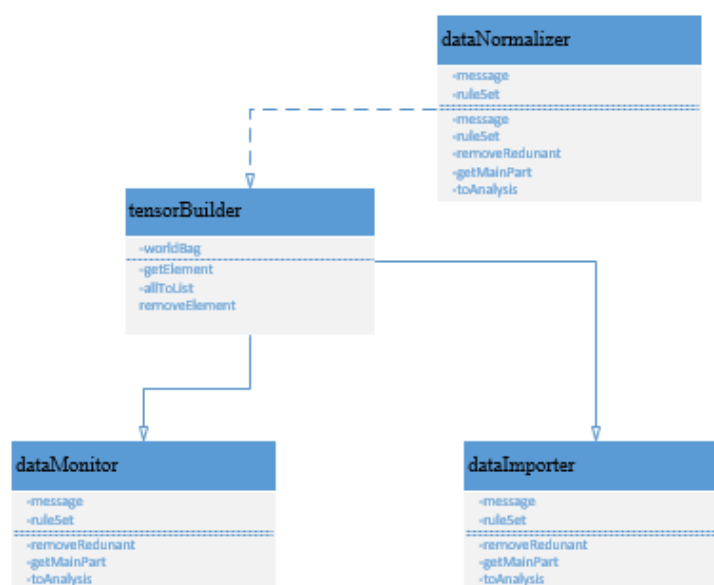


Рисунок 3.6 — Діаграма класів модулю обробки кількісних даних

3.4 Модуль для роботи із тензором

Цей модуль використовує алгоритм розкладення Такера, алгоритм скорочення розмірності та метод головних компонент, що були описані у розділі 2. Модуль повинен задовольняти наступним вимогам:

- інтегрувати зібрані дані у єдиний тензор;
- скоротити розмірність тензора за допомогою розкладання Такера та алгоритму скорочення розмірності;

3.4.3 Бібліотки, що застосовуються у модулі.

Для роботи із тензором використовується бібліотека `tensorflow` мови Python 3.6. TensorFlow є системою машинного навчання Google Brain другого покоління. У той час як еталонна реалізація працює на одиничних пристроях, TensorFlow може працювати на багатьох паралельних процесорах, як CPU, так і GPU, спираючись на архітектуру CUDA для підтримки обчислень загального призначення на графічних процесорах. TensorFlow підтримує 64-розрядні Linux, macOS, Windows, і мобільні обчислювальні платформи, включаючи Android і iOS. Обчислення TensorFlow виражаються у вигляді потоків даних через граф станів. Назва TensorFlow походить від операцій з багатовимірними масивами даних, які також називаються «тензорами». У червні 2016 року Джефф Дін з Google відзначив, що до TensorFlow зверталися 1500 репозиторіїв на GitHub, і тільки 5 з них були від Google.

Модуль не застосовує інструментів збереження даних

3.4.4 Архітектура модуля.

Модуль складається з наступних класів:

- Клас `tensorBuiler`. Формує дані, що будуть придатні до обробки у `tensorflow`.

- Клас `tensorSCMrebuilder`. Скорочує розмірність тензора використовуючи розкладення Такера та алгоритм скорочення розмірності тензору.
- Клас `tensor`. Представляє собою об'єкт тензора даних;

Діаграма класів модуля зображена на рисунку 3.7

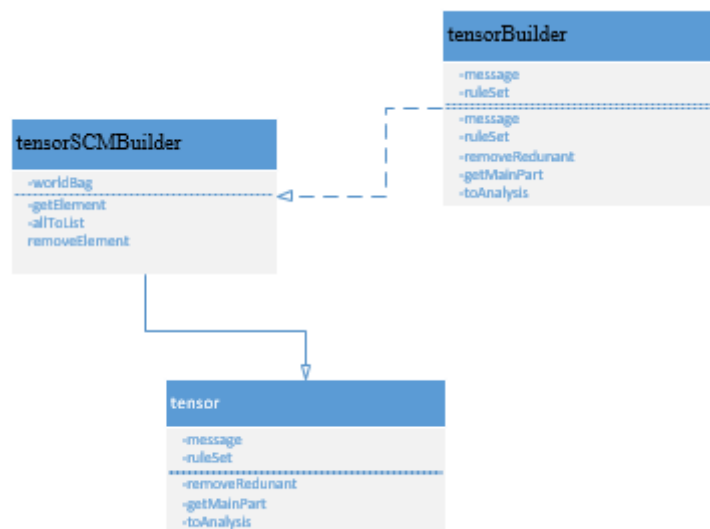


Рисунок 3.7 — Діаграма класів модулю роботи із тензором

3.5 Нейромережевий модуль

Цей модуль працює із нейронною мережею із довною короткостроковою пам'яттю, принципи роботи якої описані у розділі 2. Також модуль реалізує алгоритм навчання із зворотнім поширенням похибки у часі. Модуль повинен задовольняти наступним вимогам:

- здійснювати прогнозування часового ряду;
- мати режим навчання;
- мати можливість видалення результатів навчання;
- мати можливість до візуалізації результатів прогнозування;

3.5.1 Бібліотки, що застосовуються у модулі.

Для роботи із неймережею використовується пакет Keras. Keras — відкрита неймережева бібліотека, написана на мові Python. Вона являє собою надбудову над фреймворками DeepLearning4j, TensorFlow і Theano. Націлена на оперативну роботу з мережами глибинного навчання, при цьому спроектована так, щоб бути компактною, модульною та розширюється.

Модуль використовує СУБД SQLite3 для збереження результатів прогнозування. Структура неймережі зображена на рисунку 3.8

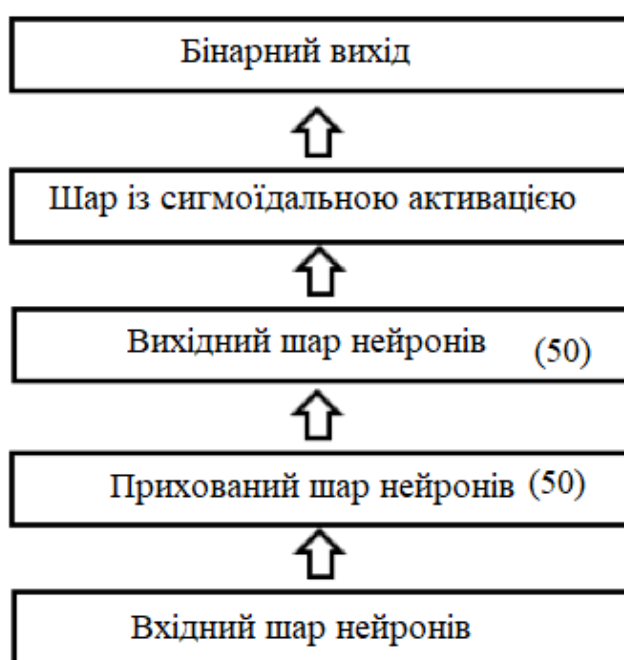


Рисунок 3.8 — Структура неймережі

3.6 Тестування системи

Протестуємо розроблену систему згідно вимог, поставлених у розділі 2. Перевіримо основні вимоги за допомогою тестових сценаріїв.

У таблиці 3.1 наведено тест-кейс активації ліцензії

Таблиця 3.1 — Тест-кейс активації ліцензії.

Назва:	Тест активації ліцензії		
Функція:	Активация ліцензії		
Дія	Очікуваний результат	Результат тесту: пройшов не пройшов заблокований	
Передумова:			
Запустіть неактивоване застосування	Запущене готове до роботи застосування	Пройшов	
Перейдіть на вкладку “активация”	Відкрита активна вкладка активация. Запуск тесту.	Пройшов	
Кроки тестування:			
Введіть невалідний ключ ліцензії	З'являється нове вікно із повідомленням про невалідність ключу	Пройшов	
Введіть валідний ключ ліцензії	З'являється вікно із повідомленням про успішну активацию	Пройшов	
Постумова:			
Стають активними вікна роботи із нейромережею та модулем збору новин	З'являється можливість використовувати модуль збору новин	Пройшов	
Підгружається база даних	З'являється можливість робити прогнози	Пройшов	

У таблиці 3.2 наведено тест-кейс збору фінансових новин

Таблиця 3.2 — Тест-кейс збору фінансових новин.

Назва:	Тест збору фінансових новин		
Функція:	Збір фінансових новин		
Дія	Очікуваний результат	Результат тесту: пройшов не пройшов заблокований	
Передумова:			
Заповніть або доповніть список сайтів-постачальників фінансових новин	Заповнений та активний список сатів-постачальників фінансових новин.	Пройшов	
Натисніть кнопку “Моніторинг”	Запуск тесту.	Пройшов	
Кроки тестування:			
Перейти в режим моніторингу	Відкрита вкладка моніторингу із періодично додаваними адресами сторінок	Пройшов	
Ввести у поле “Slack” webkook-токен чату	Переадресація нових веб-сторінок у слек-чат	Пройшов	
Постумова:			
Залишити систему у стані моніторингу	Продовження моніторингу до відключення системи, регулярні сповіщення у слек-чаті	Пройшов	
Ввести назву компанії у поле пошуку	Фільтрація сторінок згідно с критерієм	Пройшов	

У таблиці 3.3 наведено тест-кейс перевірки виявлення позитивної тональності

Таблиця 3.3 — Тест-кейс перевірки тональності

Назва:	Тест перевірки позитивної тональності		
Функція:	Аналіз тональності повідомлення		
Дія	Очікуваний результат	Результат тесту: пройшов не пройшов заблокований	
Передумова:			
Виконати дії описані у попередньому тест кейсі	Список повідомлень або новин у текстовому вигляді у таблиці	Пройшов	
Перейти до вкладки “Аналіз тональності”	Запуск тесту.	Пройшов	
Кроки тестування:			
Натиснути кнопку “Перевірка тональності”	Відкрите вікно “Перевірка тональності”	Пройшов	
Ввести у поле вводу щось позтивне англійською мовою, наприклад “So good so happy))”	Позитивне значення у полі тональності.	Пройшов	
Постумова:			
Закриття вікна перевірки тональності	Повернення у звичаний режим роботи	Пройшов	
----	---	Пройшов	

У таблиці 3.4 наведено тест-кейс перевірки виявлення негативної тональності

Таблиця 3.4 — Тест-кейс перевірки тональності

Назва:	Тест перевірки негативної тональності		
Функція:	Визначення тональності повідомлення		
Дія Запустіть застосування	Очікуваний результат	Результат тесту: пройшов не пройшов заблокований	
Передумова:			
Виконати дії описані у попередньому тест кейсі	Список повідомлень або новин у текстовому вигляді у таблиці	Пройшов	
Перейти до вкладки “Аналіз тональності”	Запуск тесту.	Пройшов	
Кроки тестування:			
Натиснути кнопку “Перевірка тональності”	Відкрите вікно “Перевірка тональності”	Пройшов	
Ввести щось негативне англійською мовою, наприклад “Bad and ugly”	Позитивне значення у полі “Тональність”.	Пройшов	
Постумова:			
Закриття вікна перевірки тональності	Повернення у звичайний режим роботи	Пройшов	
--	--	Пройшов	

У таблиці 3.5 наведено тест-кейс ручного імпорту даних

Таблиця 3.5 — Тест-кейс ручного імпорту даних

Назва:	Тест ручного імпорту кількісних даних
--------	---------------------------------------

Продовження таблиці 3.5

Продовжити наслідок			
Функція:	Імпорт кількісних даних		
Дія	Очікуваний результат	Результат тесту: пройшов не пройшов заблокований	
Передумова:			
Запустіть застосування, активуйте його.	Готове до роботи активоване застосування	Пройшов	
Натисніть кнопку “Ручний імпорт”	Запуск тесту.	Пройшов	
Кроки тестування:			
Виберіть файл який не відповідає формату введення кількісних даних	З'явиться вікно “Невірний формат”	Пройшов	
Виберіть файл який не відповідає формату введення кількісних даних	З'явиться вікно “Імпорт виконано”	Пройшов	
Постумова:			
Застосуйте дані ручного імпорту для навчання нейромережі	Нейромережу навчену	Пройшов	
Видаліть дані ручного імпорту	Дані видалено	Пройшов	

У таблиці 3.6 наведено тест-кейс навчання нейронної мережі

Таблиця 3.6 — Тест-кейс навчання нейронної мережі

Назва:	Тест навчання нейроноої мережі
Функція:	Навчання нейроннох мережі

Продовження таблиці 3.6

Дія Запустіть застосування	Очікуваний результат	Результат тесту: пройшов не пройшов заблокований
Передумова:		
Запустіть застосування, активуйте його.	Готове до роботи активоване застосування	Пройшов
Натисніть кнопку “Навчання”	Запуск тесту.	Пройшов
Кроки тестування:		
Натисніть кнопку “Cuba”	Відкриється вікно з вибором активу та дат навчальної виборки	Пройшов
Натисніть кнопку навчання	Деякий час програма буде неактивною, потім з'явиться вікно “Навчено!”	Пройшов
Постумова:		
Зробіть прогнозування на вже відомий день у минулому	Прогнозування виконане	Пройшов
Зробіть прогнозування на тестовій виборці	Прогнозування виконане	Пройшов

У таблиці 3.7 наведено тест-кейс прогнозування фінансового активу

Таблиця 3.7 — Тест-кейс прогнозування фінансового активу

Назва:	Тест прогнозування тренду фінансового активу
Функція:	Прогнозування

Продовження таблиці 3.7

Дія Запустіть застосування	Очікуваний результат	Результат тесту: Пройшов не пройшов заблокований
Передумова:		
Запустіть застосування, активуйте його.	Готове до роботи активоване застосування	Пройшов
Натисніть кнопку “Навчання”	Запуск тесту.	Пройшов
Кроки тестування:		
Натисніть кнопку “Cuba”	Відкриється вікно з вибором активу та дат навчальної виборки	Пройшов
Натисніть кнопку навчання	Деякий час програма буде неактивною, потім з'явиться вікно “Навчено!”	Пройшов
Постумова:		
Проведіть прогнозування для N майбутніх днів	Вірне прогнозування у 55% випадків	Пройшов
Проведіть прогнозування для відносно дат навчання минулого дня	З'явиться вікно “Не можна робити прогнози на минулі дати”	Пройшов

Висновки до розділу

Результатом розробки стала система прогнозування фінансових ринків із використанням рекурентних нейромереж, що включає у себе модулі аналізу повідомлень у соціальних мережах, модуль обробки фінансових новин, модуль

обробки технічних параметрів, модуль скорочення тензору, та, власне, нейромережевий модуль.

При створенні системи досліджено переваги та недоліки вже існуючих на ринку систем прогнозування фінансових активів розглянутих у розділі 1. Були сформовані основні пріоритети та декомпозовано задачі, поставлені у розділі 2. Розроблена система відповідає критеріям якості, що були висунулі у розділі 2 та технічному завданню магістерської дисертації.

Результатом розробки є система прогнозування фінансових ринків. Система прогнозування фінансових активів, що була створена у магістерській дисертації має на меті прогнозування тренду фінансового активу на 1 торговельний день. Система в своїй основі використовує такі моделі як LSTM, модель аналізу тональності за допомогою словників, наївний Баєсовський класифікатор.

Варто зауважити що система загалом є сучасною дружною для користувача системою і може бути впроваджена як у бізнес-процеси підприємств що працюють із фінансовими активами, так і використовуватися індивідуальними учасниками біржово-економічних відносин. Система призначена для надання допоміжної інформації аналітику, а її мета — підвищення точності прогнозування тренду.

Досліджено переваги і недоліки розробленої системи. До переваг системи можна віднести легкість розширення, великий спектр підтримуваних форматів вхідних даних, можливість до інтеграції у більші системи за допомогою зручного API, порівняно невелику потребу у оперативній і постійній пам'яті електронно-обчислювальної машини а також кросплатформеність, оскільки система використовує кросплатформені технології, мови програмування та бібліотеки. До недоліків системи можна віднести порівняно довгий час навчання, обмеженість у роботі із даними що не відповідають вхідному формату, відсутність до можливостей конфігурації та порівняно високі вимоги до процесору електронно-обчислювальної машини.

4 МОДЕЛЮВАННЯ ТА АНАЛІЗ ОТРИМАНИХ РЕЗУЛЬТАТІВ

4.1 Моделювання

Розроблена система була навчена та протестована на двох наборах даних — 8 популярних акцій NASDAQ та 5 популярних акцій гонконгської біржі. Кількісні дані (ціна відкриття, закриття, мінімальна та максимальна ціна впродовж дня, коефіцієнти волатильності та базові індекси) були взяті з ресурсу Wind.

Обробка новин та даних з соціальних мереж була проведена у дослідженні[33], за допомогою ресурсу Wind та агрегатора Guba. Обробка подій та тексту новин та приведення їх до даних, що будуть подані на вхід нейромережі є цікавою темою, що була освітлена у дослідженні[34]. Це дослідження не ставить перед собою мети покращити ці методи і використовує готові дані у публічному доступі.

Отримані дані були поділені на 3 вибірки — навчальну, тестову та перевірочну. Була проведена обробка вхідних даних.

По-перше була встановлена межа чутливості зміни тренду у 2%. За час, що відповідає навчальній вибірці 53% торгівельних сесій мали спадаючий тренд, 45 — зростаючий і 2% не показали зміни тренду. Межа чутливості у 2% потрібна для того, щоб відкинути незначні зміни курсу, що не залежали від вагомих інфоповодів. Також відкинуті дані коли зміна курсу акцій компаній не відрізнялася від основних індексів (Dow Jones та S&P500) на більше ніж 0.1%.

4.2 Результатаи моделювання.

Оцінюючи результати прогнозування, будемо користуватися двома метриками: відсотком правильного прогнозу та коефіцієнтом кореляції Метьюза. Коефіцієнт кореляції Метьюза(K) враховує значення матриці спряженості та обчислюється за наступною формулою(4.1):

$$K = \frac{IP * IN - XP * XH}{\sqrt{(IP + XP) * (IP + XH) * (IN + XP) * (IN + XH)}} \quad (4.1)$$

де IP — Істинопозитині прогнози, IN — Істинонегативні прогнози, XP — Хибнопозитивні прогнози, XH — Хибнонегативні прогнози.

4.2.1 Оцінка точності прогнозування

Результати моделювання наведені у Таблиці 4.1.

Таблиця 4.1 — Результати експерименту

	Вибірка 1 (NASDAQ)		Вибірка 2 (Гонконг)	
Метод	Точність(%)	K	Точність(%)	K
LSTM + вектор ознак	59.63%	0.168	59.52%	0.171
LSTM (тільки часовий ряд)	54.67%	0.013	55.03%	0.08
LSTM + тензор	57.70%	0.103	57.53%	0.091
LSTM + скорочений тензор	61.03%	0.189	60.81%	0.203
RNN (тільки часовий ряд)	53.27%	0.01	54.18%	0.05
RNN + тензор	56.53%	0.093	56.02%	0.094
RNN + скорочений тензор	56.63%	0.167	56.08%	0.173

Із результатів, наведених у таблиці можна зробити висновок, що мета дослідження була досягнута. Точність прогнозування була підвищена на 1.13% у порівнянні з прогнозуванням із урахуванням єдиного вектору ознак. У порівнянні із прогнозуванням без урахування зовнішніх факторів точність підвищена на 3.31%.

У таблиці приведено три виміри. Перший — прогнозування часового ряду за допомогою LSTM — нейромережі без урахування аналізу новин та соцмереж. Другий рядок — прогнозування з урахуванням аналізу новин та соцмереж. У третьому рядку додатково застосовується алгоритм зменшення розмірності тензору. Окрім того приведені виміри для класичної нейромережі RNN.

Результати експерименту показали, що прогнозування часового ряду без урахування аналізу новин та соціальних мереж дає результат 54-57%. Цей результат є вищим за випадкове прогнозування, але все ж близький до звичайного підкидання монети. Урахування аналізу соціальних мереж, новин та кількісних ознак змінює ситуацію. Точність прогнозування на рівні 59-60% є значним покращенням у порівнянні зі звичайним прогнозуванням часового ряду. Застосування алгоритму зменшення розмірності тензору покращує точність прогнозування на 1-1.3%, що також добре. Результат у 60.81% вірного прогнозу можна вважати позитивним результатом дослідження. Цей результат перевершує результат дослідження[15]. Важко виявити який саме результат може бути максимальним. Зрозуміло, що якщо гіпотетична система прогнозування давала би результат 80-100%, вона би докорінно змінила стан речей на фондових ринках, адже її користувач міг би заробити дуже великі гроші застосовуючи її. Але саме застосування такої гіпотетичної системи призведе до змін у торгівлі і поведінці учасників торгівлі, що нівелює переваги її користувача.

Системи із 50-55% вірним прогнозуванням взагалі майже не дають переваги користувачу. Фактично, процент вірного прогнозу більший за 50, але це може бути зумовлено лише застосуванням на вибірці, яка схожа на навчальні дані і застосовуючи систему у інших умовах можна не отримати результату. Результат у 61% є впевнено-позитивним. Безумовно, система не може замінити повноцінного експерта та аналітика через власну обмеженість, але застосування нового підходу та

алгоритму зменшення розмірності забезпечує більшу точність прогнозування для рекурентних нейромереж.

4.2.2. Оцінка швидкості навчання

Дослідимо швидкість навчання нейромережі для методів, наведених у пункті 4.2.1. для однакової кількості епох. Результати оцінки швидкості навчання наведені у Таблиці 4.2.

Таблиця 4.2 — Результати оцінки швидкості навчання

Метод	Час (с)
LSTM + вектор ознак	7851
LSTM (тільки часовий ряд)	2858
LSTM + тензор	16957
LSTM + скорочений тензор	14755
RNN (тільки часовий ряд)	5853
RNN + тензор	9455
RNN + скорочений тензор	6852

З таблиці можна зробити висновок що розроблена система програє у швидкості навчання усім системам-аналогам, окрім системи із нескорочений тензором. Це пояснюється більш складною структурою нейронної мережі і більшою кількост хідних даних. Цей недолік нівелюється кращим відсотком вірних прогнозів.

4.2.3. Оцінка потенційного прибутку від застосування системи

Позитивним результатом прогнозування є такий результат, що є більшим за 50%, оскільки система яка робить випадкове прогнозування буде мати саме такий результат. Будемо називати додатковим відсотком системи той відсоток, який

показує насільки результат більший за систему випадкового прогнозування. Оцінка додаткових відсотків систем прогнозування наведена у Таблиці 4.2.

Таблиця 4.2 — Додаткові відсотки систем прогнозування

	Вибірка 1 (NASDAQ)	Вибірка 2 (Гонконг)
Метод	Дод. відсоток(%)	Дод. відсоток (%)
LSTM + вектор ознак	9.63%	9.52%
LSTM (тільки часовий ряд)	4.67%	5.03%
LSTM + тензор	7.70%	7.53%
LSTM + скорочений тензор	10.03%	10.81%
RNN (тільки часовий ряд)	3.27%	4.18%
RNN + тензор	6.53%	6.02%
RNN + скорочений тензор	6.63%	6.08%

З таблиці 4.2 можна зробити висновок, що всі системи можуть приосити прибуток користувачеві. Проте існує потреба у розробці математичної моделі порівняльної оцінки прибутку, який можна отримати застосовуючи ту чи іншу систему.

Розробимо просту математичну модель оцінки прибутку, який може принести система своєму користувачу. Безумовно, існують методи оцінки ризиків, економічного планування та сформовані біржеві стратегії, проте при порівнянні систем доцільно застосувати найпростіші з них, оскільки нас цікавить саме порівняння систем, а не конкретне кількісне значення.

Нехай N — початковий капітал. M — кількість ітерацій прогнозування, k — коефіцієнт розміру інвестиції, який показує яка частка початкового капіталу припадає на кожну інвестицію. Припустимо що кожна інвестиція подвоюється якщо результат прогнозування вірний і віднімається якщо він помилковий. Це спрощена модель, але вона дозволяє адекватно порівнювати між собою системи.

Розглянемо дві базові біржеві стратегії, за допомогою яких будемо порівнювати системи. Стратегію стабільної торгівлі та стратегію ризикованої торгівлі.

Стратегія стабільної торгівлі полягає в тому, що вона дозволяє користувачу гарантовано зробити деяку кількість інвестицій, незалежно від результатів попередніх, навіть якщо всі вони були невдалими. Для коректності експерименту, будемо вважати що кількість ітерацій є статистично достовірною, тобто достатньо велико, щоб розрахований процент вірного прогнозування системи було досягнуто. Кожної ітерації користувач робить інвестицію розміром N/k , при чьому $k = M$. Таким чином загальний прибуток можна описати формулою 4.2:

$$V = \left(\frac{N}{k} * M * \frac{P}{100} * 2 \right) - \left(\frac{N}{k} * M * \frac{(100-P)}{100} * 2 \right) \quad (4.2)$$

де P – процент вірного прогнозу, розрахований у пункті 5.1

Стратегія ризикованої торгівлі відрізняється від стратегії стабільної торгівлі тим, що при збереженні розміру інвестицій в два рази збільшується кількість ітерацій. Безумовно, така стратегія не гарантує що користувач зможе зробити всі інвестиції. Ризики при такій стратегії значно більші, проте і прибуток також більший. Таким чином загальний прибуток для стратегії ризикованої торгівлі можна описати формулою 4.3:

$$V = \left(\frac{N}{k} * 2 * M * \frac{P}{100} * 2 \right) - \left(\frac{N}{k} * 2 * M * \frac{(100-P)}{100} * 2 \right) \quad (4.3)$$

де P – процент вірного прогнозу, розрахований у пункті 5.1

Результати застосування стратегії стабільної торгівлі наведені у Таблиці 4.2.

Таблиця 4.2 — Результати застосування стратегії стабільної торгівлі

	Вибірка 1 (NAS-DAQ)	Вибірка 2 (Гонконг)
Метод	Прибуток(%)	Прибуток (%)
LSTM + вектор ознак	19.26	19.04
LSTM (тільки часовий ряд)	9.34	10.06
LSTM + тензор	15.4	15.06
LSTM + скорочений тензор	20.06	21.62
RNN (тільки часовий ряд)	6.54	8.36
RNN + тензор	13.06	12.04
RNN + скорочений тензор	10.26	10.16

З таблиці видно, що розроблена система є оптимальною за критерієм отриманого прибутку серед представлених систем, проте істотним недоліком є порівняно довгий час навчання. Візуалізацію результатів застосування стратегії стабільної торгівлі наведено на рисунку 4.1

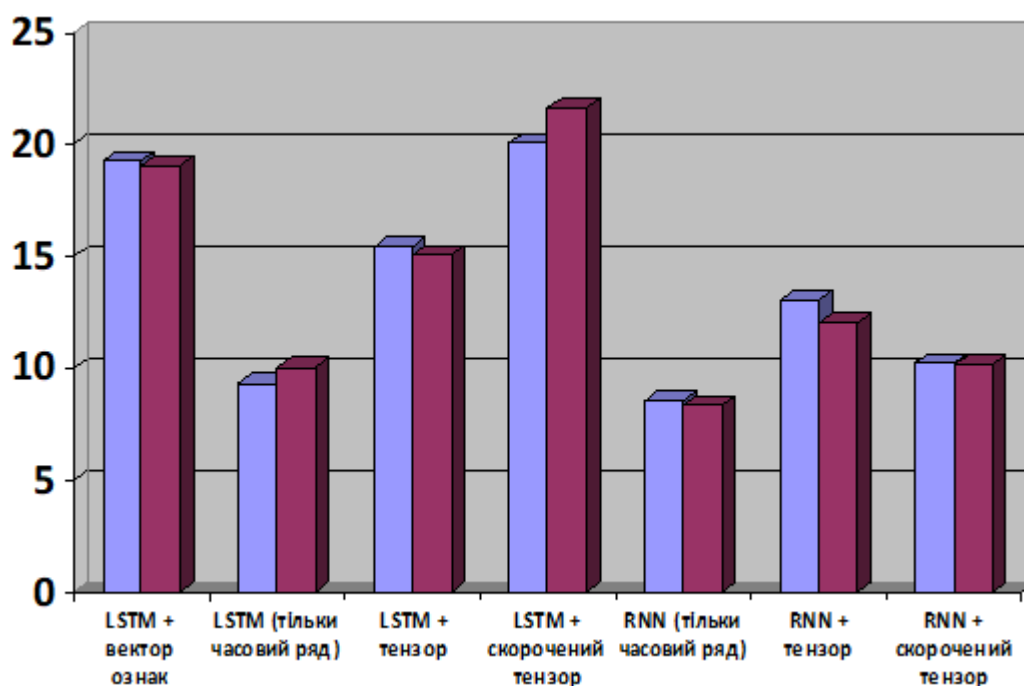


Рисунок 4.1 — Візуалізація результатів стратегії стабільної торгівлі

Результати застосування стратегії ризикованої торгівлі наведені у Таблиці 4.3.

Таблиця 4.3 — Результати застосування стратегії ризикованої торгівлі

	Вибірка 1 (NASDAQ)	Вибірка 2 (Гонконг)
Метод	Прибуток(%)	Прибуток (%)
LSTM + вектор ознак	42.22	41.70
LSTM (тільки часовий ряд)	19.55	21.1
LSTM + тензор	33.17	32.38
LSTM + скорочений тензор	44.14	47.91
RNN (тільки часовий ряд)	13.50	17.41
RNN + тензор	27.82	25.52
RNN + скорочений тензор	32.22	32.38

З таблиці видно, що розроблена система є найбільш прийнятною для застосування стратегії ризикованої торгівлі. Візуалізацію результатів застосування стратегії ризикованої торгівлі наведено на рисунку 4.2

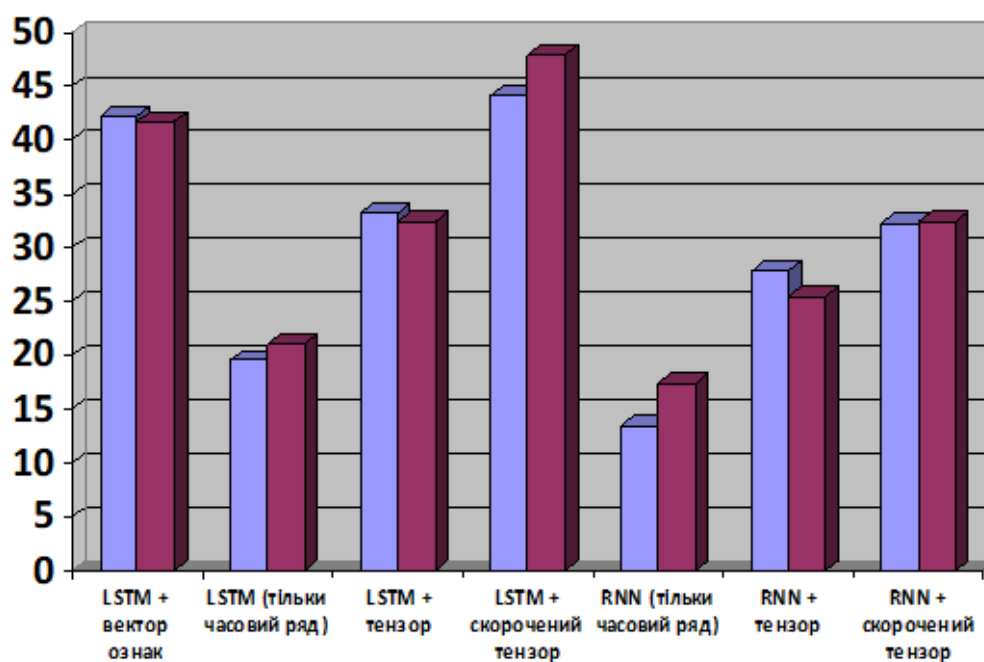


Рисунок 4.2 — Візуалізація результатів стратегії ризикованої торгівлі

Таким чином, було доведено переваги розробленої системи над системами-аналогами, та системами, що застосовують інші архітектури нейронних мереж та способи подання даних. Розплатою за кращий процент прогнозування є значно довший час навчання.

Висновки до розділу

Розроблену систему для прогнозування курсів фінансових активів було навчено та протестовано на реальних даних. Було виконано моделювання та зроблено прогнозування зміни тренду на наступний торгівельний день.

Була досягнута мета дослідження — точність прогнозування підвищена на 1.13% у порівнянні з прогнозуванням із урахуванням єдиного вектору ознак. Експеримент показав, що запропонований метод може підвищити точність прогнозування. Результат був досягнутий за допомогою комбінованого підходу до факторів що впливають на курси акцій що базується на тензорному підході. Також дослідження показало що LSTM-мережі загалом краще справляються із прогнозуванням часових рядів ніж звичайні RNN нейромережі. Подальші дослідження доцільно спрямувати на оптимізацію алгоритмів обробки текстів новин, виявлення настроїв інвесторів та покращення алгоритмів навчання нейронних мереж.

Результати експерименту показали, що прогнозування часового ряду без урахування аналізу новин та соціальних мереж дає результат 54-57%. Цей результат є вищим за випадкове прогнозування, але все ж близький до звичайного підкидання монети. Урахування аналізу соціальних мереж, новин та кількісних ознак змінює ситуацію. Точність прогнозування на рівні 59-60% є значним покращенням у порівнянні зі звичайним прогнозуванням часового ряду. Застосування алгоритму зменшення розмірності тензору покращує точність прогнозування на 1-1.3%, що також добре. Результат у 60.81% вірного прогнозу можна вважати позитивним результатом дослідження.

РОЗРОБКА СТАРТАП-ПРОЕКТУ

На сьогоднішній день все більше стартапів перетворюють ідеї пов'язані із нейромережевими технологіями у працюючі бізнес моделі. Інтерес інвесторів до стартапів, що працюють із нейронними мережами зростає з кожним роком. Все більше і більше з'являється прикладів успішного застосування штучного інтелекту у найрізноманітніших галузях діяльності людини, значно зростає кількість корпорацій що впроваджують системи прогнозування із застосуванням нейромереж у свою операційну діяльність.

Окрім цього, нейронні мережі знайшли широке застосування в прогнозуванні фінансових ринків. Нова і перспективна технологія швидко привернула увагу венчурних інвесторів. У цьому розділі розглянута розробка стартап проекту системи прогнозування фінансових ринків.

5.1 Опис ідеї проекту

Результативність застосування традиційних методів прогнозування фінансових активів, наприклад акцій, облігацій, та валюти, які вільно продаються і купуються на біржах, можна назвати недостатньою для потреб сучасного ринку. Це пов'язано із тим, що інвестиції на фондовому ринку тісно пов'язані із Інтернетом і залежні від інформаційного середовища. Для підвищення точності прогнозування доцільно застосувати таку модель, що не тільки базується на кореляціях факторів та особливостях часового ряду, а й тісно пов'язана з декількома джерелами даних.

Сучасні системи прогнозування не враховують комплексно кількісні та якісні фактори, що впливають на зміну курсів акцій або враховують у векторному вигляді, який обмежує можливості представлення вхідних даних до нейромережі. Такі системи прогнозування зазвичай мають точність 53-58% вірного прогнозу. Цієї точності замало для того, щоб використовувати такі системи як повноцінний інструмент для економічного аналізу та прогнозування. Отже, ідеєю стартап-

проекту є створення і розповсюдження системи прогнозування фінансових ринків із більшою ніж у конкурентів точністю прогнозування тренду.

Таблиця 5.1. Опис ідеї стартап-проекту

Зміст ідеї	Напрямки застосування	Вигоди для користувача
	1.Фінансова аналітика	Дешевша, порівняно з аналогами, краща точність прогнозування тренду
	2.Торівля на біржах	Можливість отримувати безпосередній прибуток за рахунок точності прогнозування

Застосування зрозумілої математичної моделі спрощує реалізацію, а доведення кращих показників прогнозування може сприяти розповсюдженню системи. Безумовно, на ринку є аналоги. Для порівняння розробленої системи проведемо аналіз потенційних техніко-економічних переваг ідеї.

Метою аналізу техніко-економічних переваг є чітке виокремлення технічних і маркетингових особливостей розробленого продукту:

- а) дослідження характеристик і властивостей розробленої системи
- б) дослідження конкурентів, товарів-аналогів, товарів-замінників та загальної ситуації на ринку де буде комерціалізуватис стартап
- в) проведення порівняльного аналізу слабких, нейтральних та сильних характеристик розробленої системи

Таблиця 5.2. Визначення сильних, слабких та нейтральних характеристик ідеї проекту.

№ п/п	Техніко-економічні характеристики ідеї	(потенційні) товари/концепції конкурентів				W (слабка сторона)	N (нейтральна сторона)	S (сильна сторона)
		Мій проєкт	Neural Builder 2015	Neural Shell	Neurall Trader			
1.	Вартість ПЗ	Низька	Висока	Висока	Висока			+

Продовження таблиці 5.2

2.	Доступність	Низький	Висока	Висока	Середній		+	
3.	Кроссплатформеність	Так	Так	Ні	Ні			+
4.	Підтримка	+	-	+	+		+	

З таблиці можна зробити висновок, що розроблена система є конкурентоспроможною.

5.2 Технологічний аудит проекту

Метою технічного аудиту є визначення переліку технологій, за допомогою яких реалізована система і їх аналіз. Визначення технологічної здійсненності ідеї проекту передбачає аналіз таких складових (таблиця 5.3):

- а) за допомогою якої технологією розроблена система згідно ідеї проекту?
- б) чи існують у відкритому доступі ці технології, чи їх потрібно додатково розробляти або купувати?
- в) чи має доступ розробник до описаних технологій?

Таблиця 5.3. Технологічна здійсненність ідеї проекту

№ п/п	Ідея проекту	Технології реалізації	Наявність технологій	Доступність технологій
		Рекурентні нейронні мережі		
		Визначення тональності тексту		
Обраною мовою програмування є Python 3.6, математичними методами – метод мішку слів, наївний класифікатор, метод головних компонент, розкладення Такера, нейронні мережі із довгою короткостроковою пам'яттю				

За результатами аналізу можна зробити висновок щодо можливості технологічної реалізації проекту. Технологічним шляхом реалізації проекту було обрано мову програмування Python в середі розробки PyCharm через їх доступність та безкоштовність.

5.3 Аналіз ринкових можливостей запуску стартап-проекту

Метою аналізу ринкових можливостей є виявлення та дослідження обмежень, можливостей та розрахунок конкретних показників реалізації розробленої системи прогнозування фінансових ринків як стартап-проекту.

Спочатку було проведено аналіз попиту: наявність попиту, обсяг, динаміка розвитку ринку (таблиця 5.4).

Таблиця 5.4 — Попередня характеристика потенційного ринку стартап-проекту

№ п/п	Показники стану ринку(найменування)	Характеристика
1	Кількість систем-конкурентів на ринку, од	4
2	Загальний обсяг продаж, грн/ум.од	234000
3	Динаміка ринку	Стагнує
4	Наявність обмежень для входу(вказати характер обмежень)	нема
5	Специфічні вимоги до стандартизації та сертифікації	нема
6	Середнє значення рентабельності в галузі(або по ринку), %	18% (відповідає середній річній ставці депозиту у гривні)

Проаналізувавши результати, можна зробити висновок що проект є придатним до інвестицій, оскільки його рентабельність перевищує відсоток депозиту.

За результатами порівняння що було наведене у таблиці 5.4 було зроблено висновок, що ринок є придатним для розповсюдження системи як стартап-проекту.

Після цього були досліджені потенційні категорії клієнтів, їх особливості та заверджено перелік вимог до системи для кожної категорії клієнтів(табл. 5.5)

Таблиця 5.5 — Характеристика потенційних клієнтів стартап-проект

№ п/п	Потреба, що формує ринок систем прогнозування	Цільова аудиторія та потенційні клієнти	Відмінності у поведінці різних потенційних цільових груп клієнтів	Вимоги споживачів до товару
1	Програмне забезпечення для прогнозування фінансових ринків	Компанії що займаються фінансовою аналітикою, трейдери, керівництво публічних корпорацій	Цільові групи клієнтів мають різні доходи і потенційний прибуток від застосування системи, різні очікування та мотивація	Зручний інтерфейс, наявність інструкції з користування, кроссплатформеність.

Після дослідження потенційних категорій клієнтів було проведено дослідження ринкового середовища: складено таблиці факторів, що сприяють реалізації розробленої системи як стартап-проекту, та факторів, що йому заважають (табл. 5.6, 5.7).

Таблиця 5.6 - Фактори загроз

№ п/п	Фактор	Зміст загрози	Можлива реакція компанії
1	Конкуренція	Вихід на ринок систем з кращими характеристиками та моделлю надання послуг	Вийти на ринок зосередивши увагу на переваги власної системи. Покращення якісних характеристик системи прогнозування. Обрати цільову аудиторію

			клієнтів і запропонувати систему із більшою точністю за меншою ціною ніж у конкурентів.
2	Зміна потреб користувачів	Клієнту потрібна буде система з більшою точністю прогнозування і новими функціями	Передбачити можливість розширення системи та підвищення точності прогнозування

Таблиця 5.7 — Фактори можливостей

№ п/п	Фактор	Зміст можливості	Можлива реакція компанії
1	Конкуренція	Відсутність аналогів на українському ринку для вітчизняного користувача	Адаптація системи до особливостей потреб на українському ринку
2	Поява альтернативних методів моделювання	З'являться нові методи моделювання, що будуть більш легкими в освоєнні та праці	Розширення можливостей, максимальне спрощення в роботі з системою для клієнтів.

Надалі було проведено дослідження пропозиції: визначили загальні характеристики конкуренції на ринку (таблиця 5.8).

Таблиця 5.8 — Ступеневий аналіз конкуренції на ринку

Особливості конкурентного середовища	В чому виражена дана характеристика	Вплив на діяльність компанії (можливі дії для підвищення конкурентоспроможності)
--------------------------------------	-------------------------------------	--

Продовження таблиці 5.8

1. Вказати тип конкуренції - монополія	На ринку присутні декілька постачальників-конкурентів, але їх товар дещо відрізняється від нашої системи.	Підтримка якості системи, безперервний розвиток, покращення, вдосконалення, оновлення та підтримка.
2. За рівнем конкурентної боротьби - міжнародний	Компанії-конкуренти з інших країн	Створити основу системи прогнозування таким чином, щоб можна було легко локалізувати її для використання у інших країнах.
3. За галузевою ознакою - внутрішньогалузева	Система може застосовуватися в одній галузі, але її різних сферах.	Постійне вдосконалення системи прогнозування, що не має прив'язки до сфери, але має до галузі
4. Конкуренція за видами товарів: - товарно-видова	Конкуренція між видами систем прогнозування, їх особливостями.	Створити систему прогнозування, враховуючи недоліки конкурентів
5. За характером конкурентних переваг - цінова	Покращення процесу створення програмного продукту, мінімізація витрат на оновлення, застосування безперервної інтеграції	Використання відкритих та дешевших технологій для побудови системи в порівнянні з системами-конкурентами, але тільки якщо ці технології відповідають необхідним критеріям якості.
6. За інтенсивністю - не марочна	Бренд присутній, але його роль незначна	Реклама, участь у конференціях, семінарах, виставках.

Було проведено аналіз конкуренції у галузі за моделлю М. Портера (табл. 5.9).

Таблиця 5.9 — Аналіз конкуренції в галузі за М. Портером

Складові аналізу	Прямі конкуренти в галузі	Потенційні конкуренти	Постачальники	Клієнти	Товари-замінники
------------------	---------------------------	-----------------------	---------------	---------	------------------

Продовження таблиці 5.9

	Навести перелік безпосередніх конкурентів	Визначити бар'єри входження в ринок	Визначити фактори сили постачальників	Визначити фактори сили клієнтів	Фактори загроз з боку тоарів-замінників
	Neural Builder	Наявність вже існуючих рішень	-	Якість системи та її підтримка оновлення	Більш відомий розробник, що підтримує свою систему
Висновки :	На даний момент немає конкурентів на українському ринку	Виход на український ринок буде легшим через відсутність конкуренції	-	Вимоги клієнтів такі, як зручний інтерфейс, якість програмного продукту	Випустити систему, що буде не гірше, ніж у конкурента, але мати кращу точність прогнозування.

За результатами порівняльного дослідження що наведене у табл. 5.9 було зроблено висновок про доцільність виходу на український та міжнародні ринки.

На основі аналізу ситуації на ринках, проведеного в табл. 5.9, а також із урахуванням характеристик розробленої системи (табл. 5.2), вимог потенційних клієнтів до товару (табл. 5.5) та особливостей ринкового середовища (таблиці 5.6, 5.7) досліджується та визначається перелік факторів конкурентоспроможності розробленої системи прогнозування фінансових ринків. Аналіз факторів конкурентоспроможності розробленої системи наведено у табл. 5.10

Таблиця 5.10 — Обґрунтування факторів конкурентоспроможності

№ п/п	Фактор конкурентоспроможності	Обґрунтування (наведення чинників, що роблять фактор для порівняння конкурентних проектів значущим)
-------	-------------------------------	---

Продовження таблиці 5.10

1	Ціна	Менша ціна приваблює нових клієнтів
2	Кросплатформність	Можливість використання продукту на будь-якій операційній системі
3	Орієнтованість на кінцевого споживача	Система орієнтована на взаємодію з потенційним клієнтом

За визначеними факторами конкурентоспроможності (табл. 5.10) проведено аналіз сильних та слабких сторін стартап-проекту (табл. 5.11).

Таблиця 5.11 — Порівняльний аналіз сильних та слабких сторін проекту

№ п/п	Фактор конкурентоспроможності	Бали 1-20	Рейтинг систем-конкурентів у порівнянні з розробленою системою прогнозування						
			-3	-2	-1	0	+1	+2	+3
1	Ціна	15					*		
2	Кросплатформність	20			*				
3	Орієнтованість на кінцевого споживача	7					*		

Кінцевим кроком маркетингового дослідження можливостей реалізації системи прогнозування фінансових ринків як стартап-проекту є побудова SWOT-аналізу (матриці сильних (Strength) та слабких (Weak) сторін, загроз (Troubles) та можливостей (Opportunities) (таблиця 5.12) на основі описаних конкурентних та маркетингових загроз та можливостей, та сильних і слабких сторін (таблиця 5.11). Список маркетингових загроз та можливостей було складено на основі дослідження факторів загроз та факторів можливостей ринкової ситуації. Маркетингові загрози та можливості є наслідками (прогнозованими результатами) впливу ринкових факторів.

Таблиця 5.12 — SWOT-аналіз стартап-проекту

Сильні сторони: Ціна, орієнтованість на кінцевого споживача, кросплатформність	Слабкі сторони: Складність розповсюджувати продукцію за кордоном.
Можливості: Відсутність конкуренції на українському ринку	Загрози: Зміна основних потреб клієнтів, при відсутності конкуренції необхідно підтримувати інтерес аудиторії до продукту

За результатами SWOT-аналізу було сформовано альтернативи ринкової стратегії (перелік заходів) для виведення розробленої системи як стартап-проекту на український та міжнародні ринки та розрахований оптимальний час їх ринкової

реалізації з урахуванням потенційних розробок конкурентів, що можуть бути виведені на ринок (див. таблицю 5.9, аналіз потенційних конкурентів). Запропоновані альтернативи були проаналізовані з точки зору часу реалізації та ймовірності отримання ресурсів (таблиця 5.13).

Таблиця 5.13 — Альтернативи ринкового впровадження стартап-проекту

№ п/п	Альтернатива (орієнтовний комплекс заходів) ринкової поведінки	Приблизна ймовірність отримання ресурсів	Приблизні строки реалізації
1	Безкоштовне розповсюдження обмеженої версії створеного програмного продукту	85%	12 місяців
2	Створення програмної системи з подальшим платним розповсюдженням (продаж платної ліцензії)	45%	12 місяців

Після аналізу було обрано альтернативу №2.

5.4 Розроблення ринкової стратегії проекту

Розроблення ринкової стратегії першим етапом передбачає визначення стратегії охоплення ринку: дослідження цільових груп потенційних клієнтів, які приведені в таблиці 5.14.

Таблиця 5.4. Вибір цільових груп потенційних споживачів

№ п/п	Опис профілю цільової групи потенційних клієнтів	Готовність споживачів сприйняти продукт	Орієнтовний попит в межах цільової групи (сегменту)	Інтенсивність конкуренції в сегменті	Простота входу у сегмент
1	Компанії що займаються фінансовою аналітикою	Готові	Необхідно	Висока	Середня
Визначена цільова група клієнтів: Компанії що займаються фінансовою аналітикою					

Дослідивши потенційні групи клієнтів було визначено цільові категорію, для яких буде пропонуватися розроблена система, та визначено стратегію охоплення ринку — стратегію диференційованого маркетингу (компанія працює з декількома сегментами ринку).

Для роботи в обраних сегментах ринку сформовано базову стратегію розвитку (таблиця 5.15).

Таблиця 5.15 — Визначення базової стратегії розвитку

№ п/п	Обрана альтернатива розвитку стартап-проекту	Стратегія охоплення ринку	Ключові конкурентоспроможні позиції відповідно до обраної альтернативи	Базова стратегія розвитку*
1		Визначити потреби сучасного ринку для кожної з груп, розробити відповідно до них стратегії приваблення покупців та маркетингової компанії	Цінова політика, універсальність продукту.	Стратегія диференціації

Наступним кроком обрано стратегію конкурентної поведінки (таблиця 5.16).

Таблиця 5.16 — Визначення базової стратегії конкурентної поведінки

№ п/п	Чи є проект першою подібною системою на ринку?	Чи буде розробник системи шукати нових споживачів, або забирати існуючих у конкурентів?	Чи буде розробник стартап-проекту копіювати основні властивості товару конкурента	Стратегія конкурентної поведінки*
1	Ні	Шукати нових	Ні	Стратегія заняття конкурентної ніші

За результатами аналізу вимог клієнтів визначених категорій до розробника стартап-проекту та до продукту (див. таблицю 5.5), а також в залежності від обраної базової стратегії розвитку (таблиця 5.15) та стратегії конкурентної поведінки (таблиця 5.6) розроблено стратегію позиціонування (таблиця 5.17), що полягає у формуванні ринкової позиції (комплексу асоціацій), за яким споживачі мають ідентифікувати торгівельну марку/проект.

Таблиця 5.17 — Визначення стратегії позиціонування

№ п/п	Вимоги до системи з боку потенційних клієнтів	Основна стратегія розвитку	Ключові конкурентоспроможні позиції власного стартап-проекту	Вибір основних асоціацій
1	Проста побудова моделі прогнозування, довгий час навчання нейромережі, проте вища точність прогнозу та інтуїтивно зрозумілий інтерфейс, докладне керівництво користувача, підтримка користувачів з боку розробника	Стратегія диференціації	Позиція на основі порівняння фірми з товарами конкурентів; Відмінні особливості споживача	Економія часу; Зручність застосування; Практичність

За результатами дослідження була сформована система рішень щодо ринкової поведінки стартап-компанії, яка визначає напрями роботи стартап-компанії українському та міжнародному ринках.

5.5 Розроблення маркетингової програми стартап-проекту

Визначено маркетингову концепцію системи прогнозування фінансових ринків із використанням рекурентних нейромереж, яку буде купувати споживач. Для цього

у таблиці 5.18 наведено результати попереднього дослідження конкурентоспроможності товару. Основна концепція продукту — письмовий опис фізичних та прогнамних характеристик системи, які сприймаються споживачем, і набору вигод, які він обіцяє певній групі споживачів.

Таблиця 5.17 — Визначення ключових переваг концепції потенційного товару

№ п/п	Потреба	Вигоди, які надає система	Основні переваги над продуктами-конкурентами
1	Швидкість отримання результату	Проста побудова моделі прогнозування, довгий час навчання нейромережі, проте вища точність прогнозу	Відсутність необхідності звертатися до компаній-аналітиків та експертів у прогнозуванні
2	Зручність застосування	Інтуїтивно зрозумілий інтерфейс, докладне керівництво користувача, підтримка користувачів з боку розробника	Використано нейромережевий підхід до прогнозування.

За результатами дослідження було сформовано трирівневу маркетингову модель системи: ідея продукту та/або послуги, його фізичні складові, особливості процесу його надання (таблиця 5.19).

1-й рівень вирішує питання щодо того, засобом вирішення якої потреби і / або проблеми буде система, яка її основна вигода

2-й рівень являє рішення того, як буде реалізована система в реальному ринковому середовищі. Рівень включає в себе якість, властивості, дизайн, упаковку, ціну.

3-й рівень визначає додаткові послуги та переваги для клієнта системи, що створюються на основі товару за задумом і товару в реальному виконанні (гарантії якості, доставка, умови оплати та ін).

Таблиця 5.19 — Опис трьох рівнів моделі товару

Таблиця 3.19. Опис трьох рівнів моделі товару			
Рівні товару	Сутність та складові		
I. Ідея стартап-проекту	Програмний продукт, який дозволяє робити прогнозування фінансових ринків із використанням рекурентних нейромереж		
II. Реальне виконання системи як стартап-проекту	Властивості/характеристики	М/Нм	Вр/Тх /Тл/Е/Ор
	1. Кросплатформеність 2. Інтуїтивна зрозумілість інтерфейсу, інтеграція з основними офісними програмами 3. Інтеграція з основними форматами біржових даних		
	Якість: стандарти, нормативи, точність прогнозування Стандартизація відповідно до ДСТУ, ISO.		
	Пакування відсутнє.		
	Марка: назва організації-розробника + назва товару		
Потенційний товар буде захищено від копіювання: патентування, сертифікати відповідності.			

Після цього визначимо цінові межі, якими необхідно керуватись при встановленні ціни на систему, яке включає аналіз ціни на товари-аналоги або товари субститути, а також аналіз рівня доходів цільової категорії клієнтів (таблиця 5.20). Аналіз проведено експертним методом.

Таблиця 5.20 — Визначення меж встановлення ціни

№ п/п	Приблизна вилка вартості товарів-замінників	Приблизна вилка вартості товарів-аналогів	Приблизний рівень доходів цільової групи клієнтів	Верхня та нижня межі вартості системи
1	800-3000\$	800-3000\$	3000\$+	200-700\$

Після цього визначимо оптимальну систему збуту, за допомогою якої буде розповсюджуватися система як сатрап-проект.(таблиця 5.21)

Таблиця 5.21 — Формування системи збуту

№ п/п	Особливості ринкової поведінки потенційних клієнтів	Функції збуту, які має виконувати постачальник системи	Глибина каналу збуту	Оптимальний канал збуту
1	Цільові клієнти – компанії, які бажають впровадити у своїй роботі сучасні засоби допоможуть прогнозувати фінансові ринки із кращою точністю	Побудова прямих контактів із потенційними клієнтами і їх підтримка. Формування попиту і стимулювання продажів. Постійна рекламна робота і ринкової інформації. Підтримка системи, кастомізація для потреб конкретного споживача.	Один (від виробника одразу споживачу)	Прямий канал збуту до клієнта, мінімізувати збутові витрати. Розвиток нових маркетингових концепцій.

Фінальною складовою маркетингової програми є визначення концепції маркетингових комунікацій (таблиця 5.22).

Таблиця 5.22 — Концепція маркетингових комунікацій

№ п/п	Особливості поведінки потенційних клієнтів	Канали комунікацій потенційних покупців	Основні особливості позиціонування проекту	Ідея рекламної політики	Концептуальні особливості рекламної політики
1	Потреба клієнтів отримати кращий прогноз фінансового ринку.	Конференції, виставки, реклама в інтернеті на тематичних сайтах.	Виставки присвячені розробці програмного забезпечення та реклама в інтернеті	Донести основну ідею, цінову політику та можливості співпраці.	Продемонструвати нове рішення та новий погляд на прогнозування фінансових ринків

Результатом дослідження стала робоча ринкова програма, що включає в себе концепції системи, збуту, просування та попереднє дослідження ціноутворення на продукт, базується на цінностях та потребах категорій покупців, конкурентні переваги системи, стан та динаміку маркетингового середовища, в межах якого впроваджено системи прогнозування фінансових ринків як стартап-проект, та відповідну обрану альтернативу ринкової поведінки.

Висновки до розділу

У цьому розділі було проведено аналіз розробленої системи прогнозування фінансових ринків у якості стартап-проекту. Варто зауважити, що проект має можливість ринкової комерціалізації, через те, що ринок систем прогнозування потребує якісного та інноваційного продукту для прогнозування фінансові ринки.

Результатом роботи є розроблений стартап-проект, план виходу на ринки програмного забезпечення та маркетингова стратегія. Розроблений стартап-проект доцільно застосувати при комерціалізації розробки.

Висока конкуренція зумовлює необхідність виваженого підходу до просування продукту. Проте кількість учасників фінансових ринків зростає з кожним днем, що зумовлює хороші перспективи для комерціалізації системи. Для впровадження ринкової реалізації проекту була обрана стратегія, яка включає розробку програмної системи із класичною моделлю ліцензування за певну плату.

Можна сказати, що подальший розвиток проекту є доцільним, оскільки кількість потенційних користувачів системи зростає кожного року.

ВИСНОВКИ ПО МАГІСТЕРСЬКИЙ ДИСЕРТАЦІЇ

В роботі було проаналізовано сьогоdnішній стан проблемної області прогнозування фінансових ринків. Було розглянуто і проаналізовано основні сучасні методи прогнозування як часових рядів, так і безпосередньо курсів фінансових активів. Виявлені переваги і недоліки кожного з методів і запропоновано метод, що поєднує переваг технічного і фундаментального аналізу з прогнозуванням за допомогою нейромереж.

Також було виконано порівняння систем, вже існуючих на ринку прогнозування фінансових активів. Виходячи з результатів порівняльного аналізу було затверджено основні пріоритети і напрямки при побудові власної системи прогнозування фінансових активів. Було описано і вибрано її архітектуру і наведений набір користувацьких вимог до системи. Після цього було обрано набір інструментів для вирішення поставлених задач.

Були розглянуті основні математичні методи прогнозування часових рядів, частковим випадком яких є зміна курсів фінансових активів. Було обрано архітектуру нейронної мережі яка найбільш доцільна для вирішення поставленої задачі — мережа з довгою короткостроковою пам'яттю. Було розглянуто можливі модифікації такої мережі.

Був запропонований тензорний підхід по подання входних даних для нейромережі. Це дозволило застосувати алгоритм скорочення розмірності тензору, тим самим підвищивши точність прогнозування.

Були досліджені алгоритми семантичної обробки природної мови і алгоритми виявлення тональності повідомлень. Було обрано підхід що базується на валентних словниках через його високі якісні характеристики та швидкість реалізації.

На основі обраних методологій і інструментів, враховуючи користувацькі вимоги було розроблено систему прогнозування фінансових ринків із використанням рекурентних нейромереж. Розроблена система виконує поставлені задачі і задовольняє поставленим вимогам.

Розроблену систему було протестовно на реальних даних і експериментально підтверджено оптимальність обраної архітектури довгої короткострокової пам'яті та методів визначення тональності тексту за критерієм точності прогнозу у відсотках.

Було розроблено стартап проект і обраховано перспективи розвитку розробленої системи як стартап-проекту та шляхи її монетизації. За результатами дослідження було зроблено висновок що розроблену систему доцільно розвивати як стартап проект.

За результатами роботи опубліковано статтю “Система прогнозування фінансових ринків із використанням рекуррентних нейромереж” у збірнику конференції “Winter InfoCom Advanced Solutions 2018”.

Подальші дослідження доцільно спрямувати на оптимізацію алгоритмів обробки текстів новин, виявлення настроїв інвесторів та покращення алгоритмів навчання нейронних мереж.

СПИСОК ПОСИЛАНЬ

1. Хайкин С. Нейронные сети, полный курс [Текст] / Саймон Хайкин. – 2-е изд., перед. – М. : Вильямс, 2008. - 1103 с. – ISBN 5-8459-0890-6
2. Dasgupta D. Designing Application-Specific Neural Networks using the Structured Genetic Algorithm [Текст] / Dasgupta D., McGregor D.: Proceedings of International Workshop on Combinations of Genetic Algorithms and Neural Networks. – 1992
3. Савчук О.В., Кривенко К.С. Інтелектуальний аналіз діагностичної інформації електро- радіокомпонентів в умовах невизначеності // Інтелектуальний аналіз інформації ІАІ- 2013 / Зб. праць. – К.: Просвіта. – 2013. – С.211-217
4. Терехов В.А., Ефимов Д.В., Тюкин И.Ю. Нейросетевые системы управления: Учеб. пособие для вузов. – М: Высшая школа. 2002. -183с.
5. Сигеру Омату. Нейроуправление и его приложения. – М.: ИПРЖР, 2000. – 272с.
6. Шеремет О. Метод опорних векторів / О. Шеремет, В. Садовой. // Мат. Мод. № 1. – 2013. – №28. – С. 13. 5. Domingos P. On the optimality of the simple Bayesian classifier under zeroone loss / P. Domingos, P. Pazzani. // Machine Learning. – 1997. – №29. – С. 103–137.
7. Зорін Ю. М. Конспект лекцій з предмету «Комп'ютерні системи штучного інтелекту» [Електронний ресурс] / Ю. М. Зорін – Режим доступу до ресурсу: <http://m.scs.kpi.ua/course/view.php?id=107>(дата звернення 08.05.2018)
8. Біла Н. І. Інформаційні системи та технології в управлінні. Теоретичні відомості і завдання до лабораторних робіт / Н. І. Біла. – С. 16–19. 8. Goodfellow I. Deep Learning / I. Goodfellow, Y. Bengio, A. Courville., 2016.
9. M. Alanyali, H. S. Moat, T. Preis, Quantifying the relationship between financial news and the stock market, Sci. Rep. 3.
10. H. Mao, S. Counts, J. Bollen, Quantifying the effects of online bullishness on international financial markets, European Central Bank Workshop on Using Big Data for Forecasting and Statistics, Frankfurt, Germany

11. LeCun Y. LeNet-5, convolutional neural networks [Электронный ресурс] / Yann LeCun – Режим доступа до ресурсу: <http://yann.lecun.com/exdb/lenet/> (дата звернення 08.05.2018)
12. Zhang W. Shift-invariant pattern recognition neural network and its optical architecture / Wei Zhang. // Proceedings of annual conference of the Japan Society of Applied Physics. – 1988. 88 11. Zhang W. Parallel distributed processing model with local space-invariant interconnections and its optical architecture / Wei Zhang. // Applied Optics. – 1990. – №29. – С. 4790–4797.
13. Subject independent facial expression recognition with robust face detection using a convolutional neural network / M. Matusugu, M. Katsuhiko, Y. Mitari, Y. Kaneda. // Neural Networks. – 2003. – №16. – С. 555–559.
14. Николенко С. И. Глубокое обучение. Погружение в мир нейронных сетей. / С. И. Николенко, А. А. Кадурин, Е. О. Архангельская., 2018. – 480 с.
14. Glorot X. Understanding the difficulty of training deep feedforward neural networks / X. Glorot, Y. Bengio. // Proceedings of the Thirteenth International Conference on Artificial Intelligence and Statistics. – 2010.
15. Матеріали курсу CS231n: Convolutional Neural Networks for Visual Recognition [Електронний ресурс] – Режим доступу до ресурсу: <https://cs231n.github.io/convolutional-networks/> (дата звернення 08.05.2018)
16. Evaluation of Pooling Operations in Convolutional Architectures for Object Recognition. / D. Scherer, A. Müller, C. Andreas, S. Behnke // Springer. – 2010. – С. 92–101
17. Graham B. Fractional Max Pooling [Електронний ресурс] / B. Graham. – 2014. – Режим доступа до ресурсу: <https://arxiv.org/abs/1412.6071> (дата звернення 08.05.2018)
17. Paul C. Tetlock. 2007. Giving content to investor sentiment: The role of media in the stock market. Journal of Finance 62, 3 (2007), 1139–1168.
18. Richard Ernest Bellman and Stuart E. Dreyfus. 1962. Applied Dynamic Programming. Rand Corporation.
19. Q. Li, Y. Chen, L. L. Jiang, P. Li, H. Chen, A tensor-based information framework for predicting the stock market, ACM Trans. Inform. Syst. (TOIS) 34 (2)

(2016) 11.

20. D. M. Cutler, J. M. Poterba, L. H. Summers, What moves stock prices, *J.Portf. Manag.* 15 (1989) 4 -12.

21. P. C. Tetlock, M. SAAR-TSECHANSKY, S. Macskassy, More than words: Quantifying language to measure firms' fundamentals, *J. Finance* 63 (3) (2008) 1437-1467.

22. R. Luss, A. D'Aspremont, Predicting abnormal returns from news using text classification, *Quantitative Finance* 15 (6) (2015) 999-1012.

23. X. Ding, Y. Zhang, T. Liu, J. Duan, Knowledge-driven event embedding for stock prediction, in: *Proceedings of the 26th International Conference on Computational Linguistics (COLING-16)*, 2016, pp. 2133-2142.

24. C.-Y. Chang, Y. Zhang, Z. Teng, Z. Bozanic, B. Ke, Measuring the information content of financial news, in: *Proceedings of the 26th International Conference on Computational Linguistics (COLING'16)*, 2016, pp. 3216-3225.

25. W.-M. Song, T. Aste, T. Di Matteo, Analysis on filtered correlation graph for information extraction, *Statistical Mechanics of Molecular Biophysics* (2008) 88.

26. F. Pozzi, T. Di Matteo, T. Aste, Spread of risk across financial markets: better to invest in the peripheries, *Scientific reports* 3.

27. R. Morales, T. Di Matteo, R. Gramatica, T. Aste, Dynamical generalized hurst exponent as a tool to monitor unstable periods in financial time series, *Physica A: Statistical Mechanics and its Applications* 391 (11) (2012) 3180– 3189.

28. N. Musmeci, T. Aste, T. di Matteo, Clustering and hierarchy of financial markets data: advantages of the dbht., *CoRR*.

29. W.-M. Song, T. Di Matteo, T. Aste, Hierarchical information clustering by means of topologically embedded graphs, *PLoS One* 7 (3) (2012) e31929.

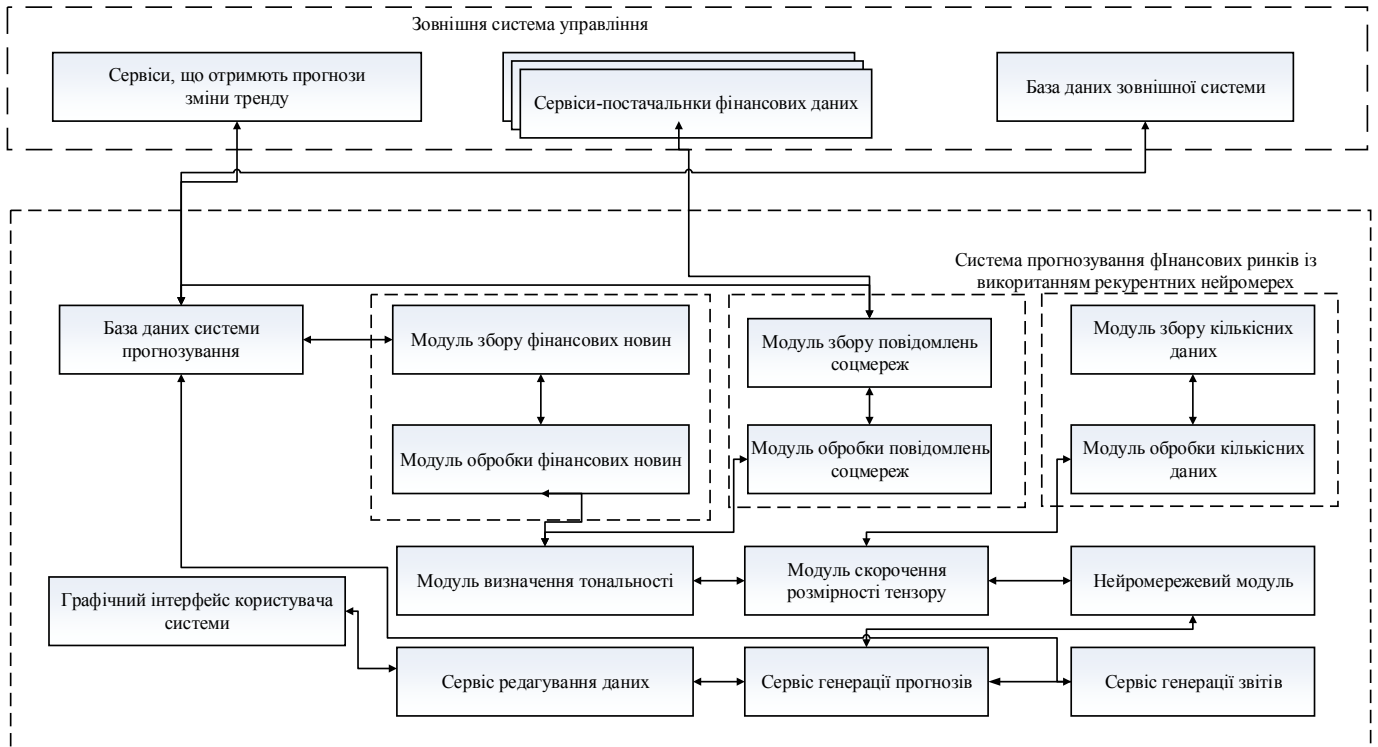
30. C. Curme, H. E. Stanley, I. Vodenska, Coupled network approach to predictability of financial market returns and news sentiments, *International Journal of Theoretical and Applied Finance* 18 (07) (2015) 1550043.

31. O. Kolchyna, T. T. P. Souza, P. Treleaven, T. Aste, Twitter sentiment analysis: Lexicon method, machine learning method and their combination, in: G. Mitra, X. Yu (Eds.), *Handbook of Sentiment Analysis in Finance*, 2016, Ch. 5.

32. J. Manfield, D. Lukacsko, T. T. P. Souza, Bull bear balance: A cluster analysis of socially informed financial volatility, in: 2017 Computing Conference, 2017, pp. 421–428. doi:10.1109/SAI.2017.8252134. 14
33. T. T. P. Souza, O. Kolchyna, P. Treleaven, T. Aste, Twitter sentiment analysis applied to finance: A case study in the retail industry, in: G. Mitra, X. Yu (Eds.), *Handbook of Sentiment Analysis in Finance*, 2016, Ch. 23.
34. I. Zheludev, R. Smith, T. Aste, When Can Social Media Lead Financial Markets?, *Scientific Reports*
35. P. C. Tetlock, Giving content to investor sentiment: The role of media in the stock market, *The Journal of Finance* 62 (3) (2007) 1139–1168.

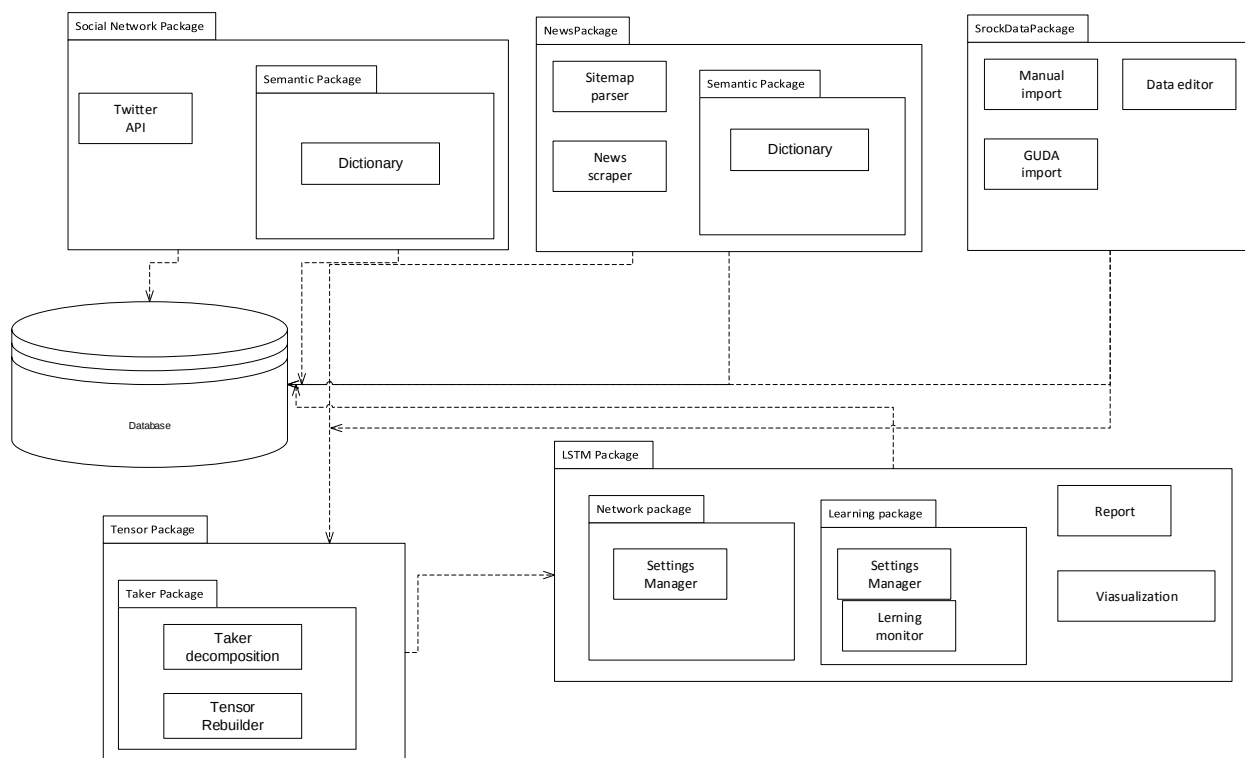
ДОДАТОК А

Структурна схема системи



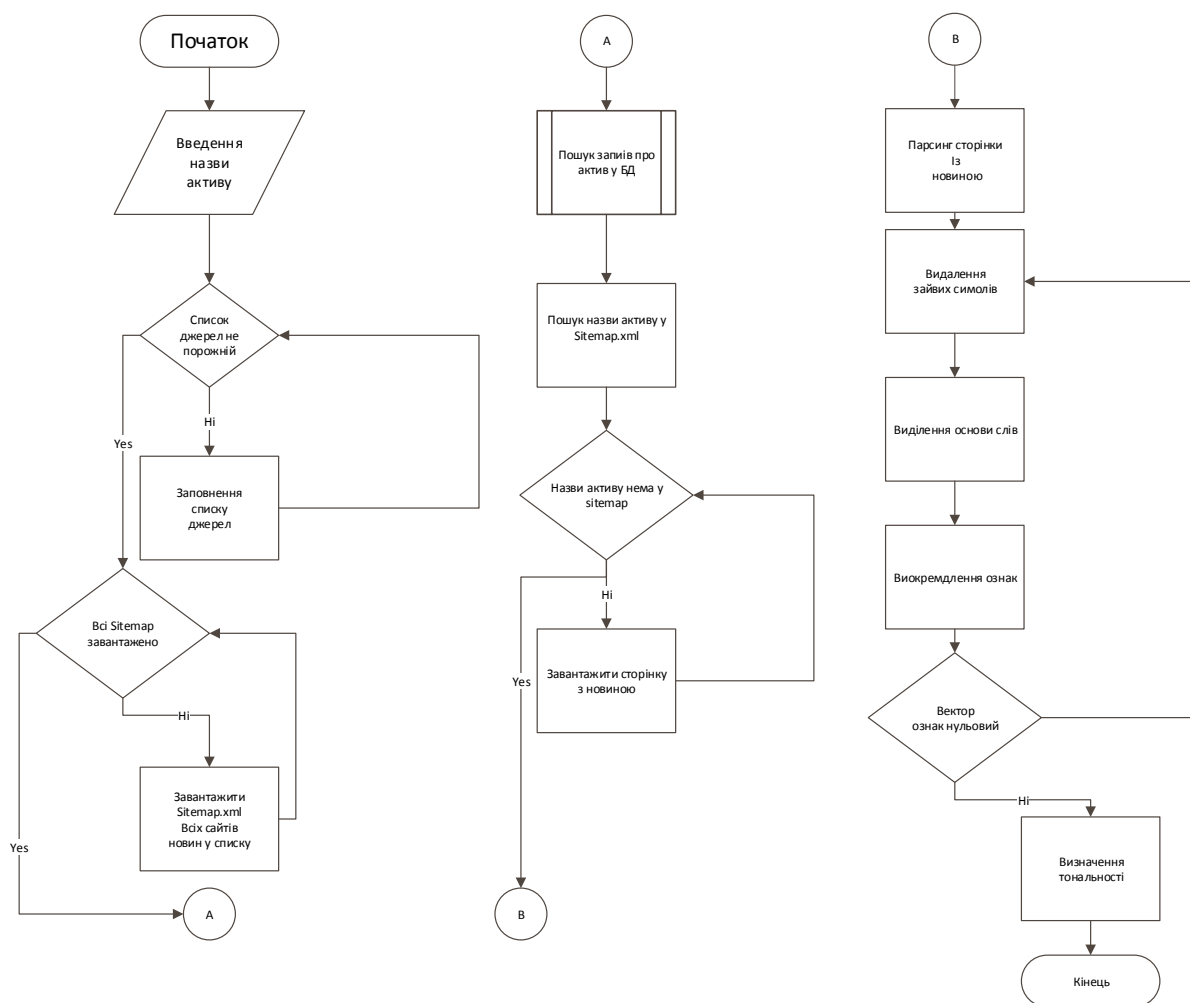
ДОДАТОК Б

Діаграма пакетів



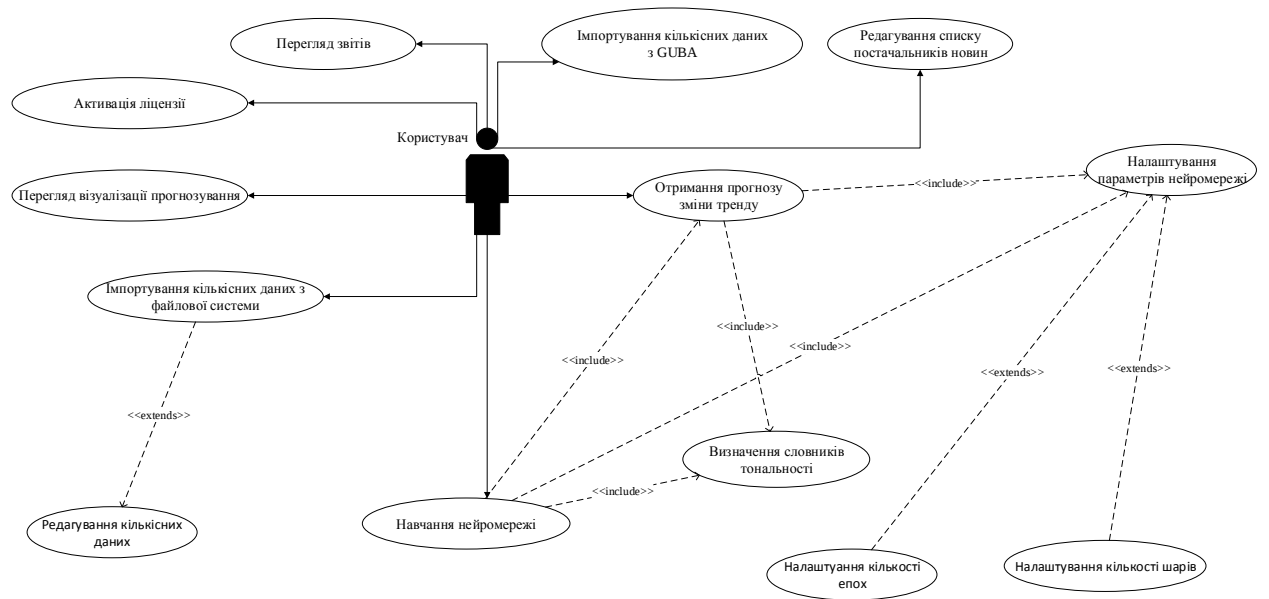
ДОДАТОК В

Алгоритм збору фінансових новин



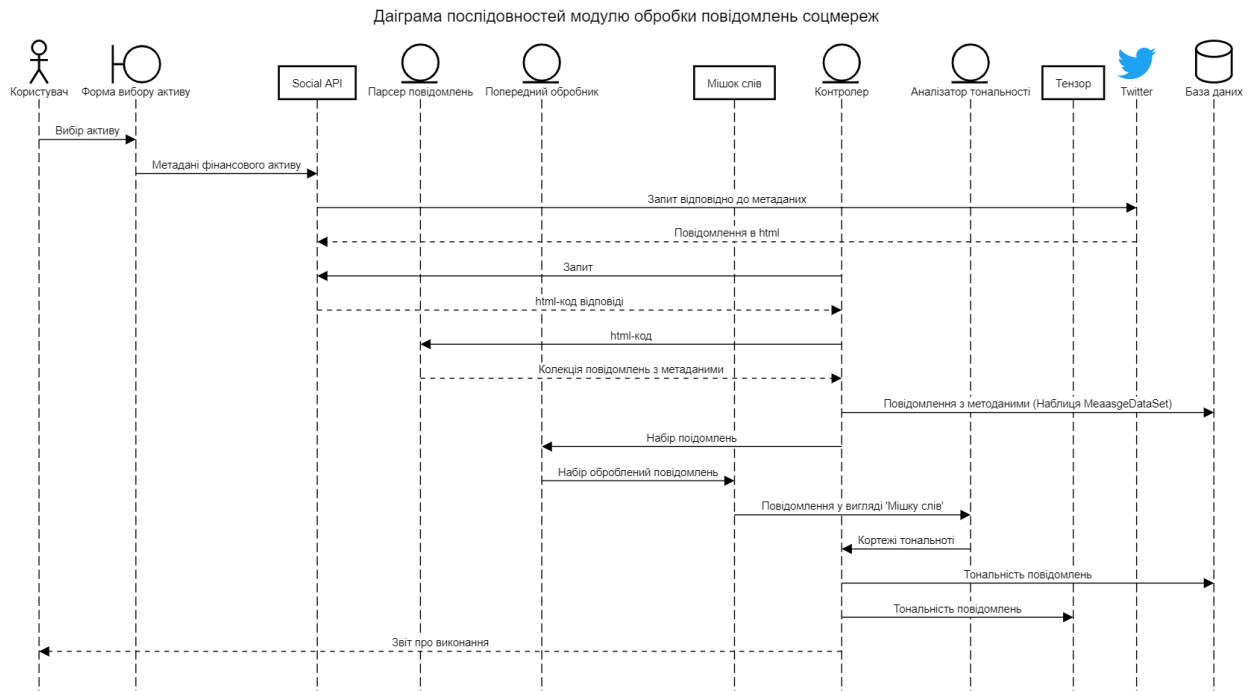
ДОДАТОК Г

Діаграма варіантів використання



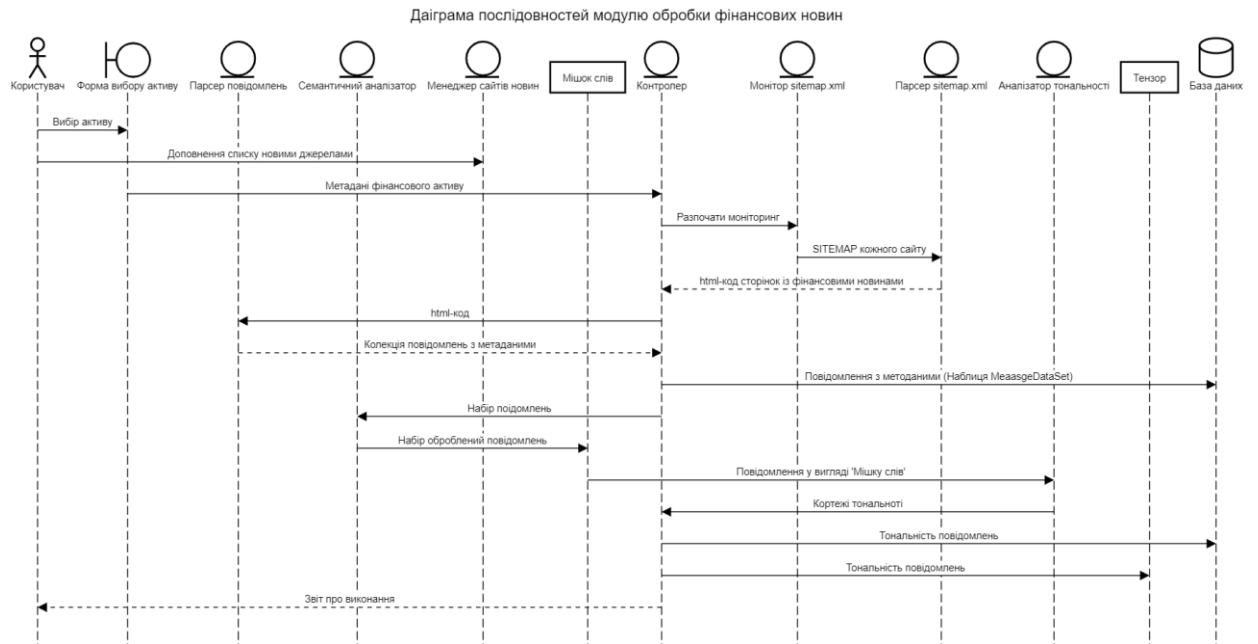
ДОДАТОК Д

Діаграма послідовностей модулю обробки повідомлень соцмереж



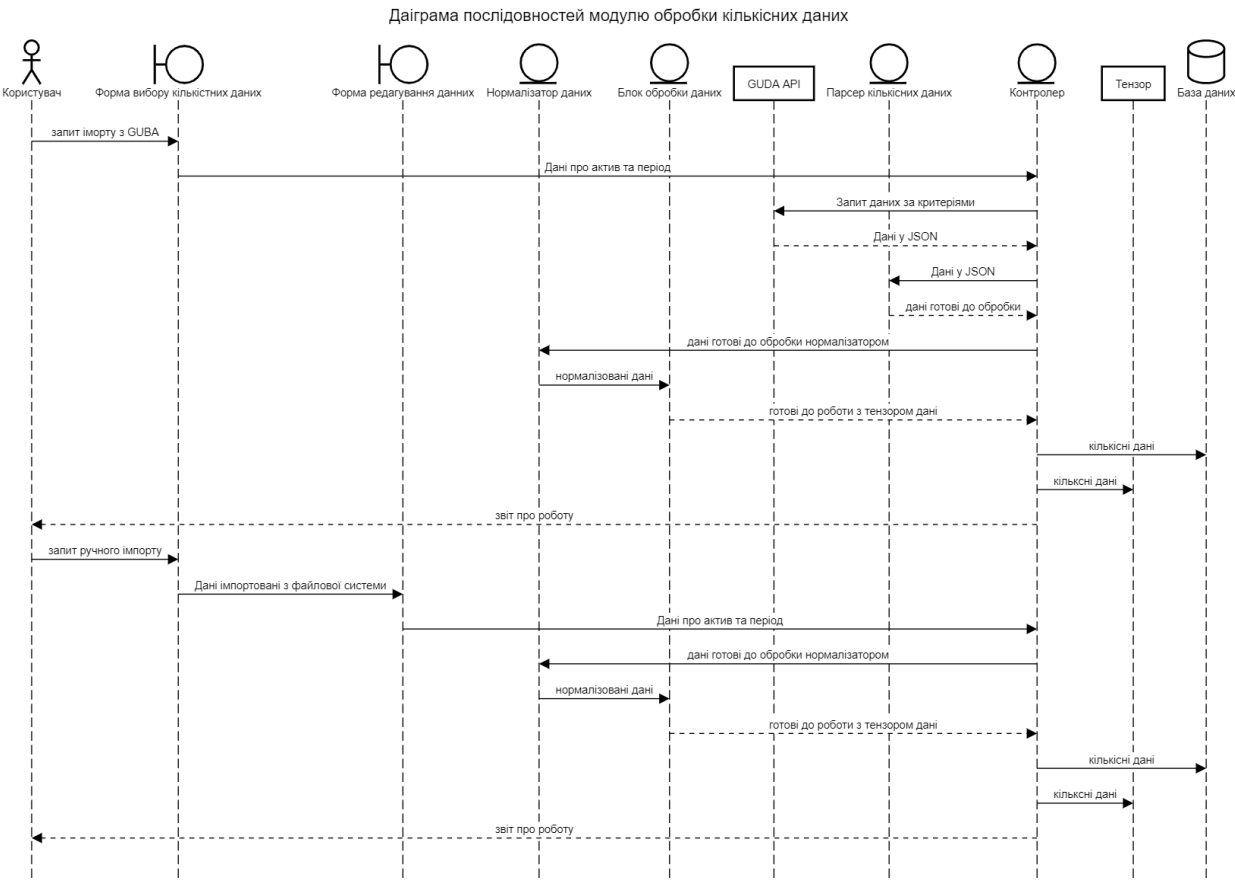
ДОДАТОК Е

Діаграма послідовностей модулю обробки фінансових новин



ДОДАТОК Ж

Діаграма послідовностей модулю обробки кількісних даних



ДОДАТОК 3

Діаграма послідовностей модулю обробки модулю прогнозування

