

ЧЕРКАСЬКИЙ НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ ІМЕНІ БОГДАНА
ХМЕЛЬНИЦЬКОГО

На правах рукопису

ЧАБАНЕНКО Дмитро Миколайович

УДК 519.216.3:519.217.2

**МАТЕМАТИЧНІ МОДЕЛІ ТА МЕТОДИ ПРОГНОЗУВАННЯ
ЧАСОВИХ РЯДІВ НА ОСНОВІ СКЛАДНИХ ЛАНЦЮГІВ
МАРКОВА**

01.05.02 — Математичне моделювання та обчислювальні методи

Дисертація на здобуття наукового ступеня кандидата
технічних наук

Науковий керівник
Соловйов Володимир Миколайович,
доктор фізико-математичних наук, професор

Черкаси — 2012

ЗМІСТ

Вступ	5
1. Аналіз проблеми прогнозування часових рядів	12
1.1. Огляд та класифікація методів прогнозування часових рядів	12
1.2. Перспективи використання трендових моделей при побудові сучасних методів прогнозування	18
1.3. Авторегресійні методи прогнозування	27
1.4. Використання нелінійних функцій, нейронних мереж та генетичних алгоритмів при побудові трендових та авторегресійних моделей.	35
1.5. Випадкові процеси та методи побудови марківських моделей при прогнозуванні.	44
Висновки до розділу 1	60
2. Методи прогнозування на основі складних ланцюгів Маркова та дискретного Фур'є-продовження	62
2.1. Метод прогнозування на основі складних ланцюгів Маркова	62
2.2. Складні ланцюги Маркова та вплив пам'яті. Порядок ланцюга Маркова	71
2.3. Ієрархія приростів часу та процедура склеювання	77
2.4. Прогнозування тестових сигналів	86
2.5. Дискретне Фур'є-продовження низькочастотної складової рядів економічної динаміки.	90
Висновки до розділу 2	95
3. Побудова інформаційної системи прогнозування	96

3.1. Аналіз середовищ розробки програмного забезпечення для прогнозування	96
3.2. Паралельні алгоритми в методі ланцюгів Маркова та програмні засоби їх реалізації	99
3.3. Проектування і розробка веб-системи розподіленого прогнозування	107
Висновки до розділу 3	115
4. Практичне використання та тестування методики складних ланцюгів Маркова	116
4.1. Обґрунтування та експериментальна апробація значень параметрів методу прогнозування, оснований на складних ланцюгах Маркова	116
4.2. Тестування прогнозу фінансового часового ряду методом складних ланцюгів Маркова	128
4.3. Багатосценарний підхід	146
4.4. Обґрунтування та експериментальне визначення параметрів методу дискретного Фур'є-продовження	148
4.5. Дискретне Фур'є-продовження з двочастинним наближенням	154
Висновки до розділу 4	158
Висновки	160
Список використаних джерел	161
Додаток 1. Результати точності представлення, квантифікації та наступного відновлення часового ряду	171
Додаток 2. Результати прогнозування часового ряду ARIMA-моделями	175
Додаток 3. Результати прогнозування часового ряду нейромережними моделями	176

Додаток 4. Результати прогнозування часового ряду методом
дискретного Фур'є-продовження 181

ВСТУП

Прогнозування часових рядів є надзвичайно актуальною задачею в дослідженні фінансово-економічних та інших складних систем. Специфіка реальних часових рядів полягає в наявності стилізованих фактів, таких як негаусівський розподіл приростів ряду, кластеризація волатильності та інші. Ці властивості є результатом складності структури досліджуваних систем та не завжди адекватно відтворюються сучасними методами прогнозування.

Фондовий ринок виконує важливу роль оптимального перерозподілу капіталу в розвинутих економіках світу. Успішне залучення коштів від розміщення цінних паперів на фондовому ринку, а також успішне інвестування можливе при глибокому розумінні законів функціонування фондового ринку. Значний доробок у дослідження фондового ринку внесли У.Ф. Шарп, Г. Марковіц, Дж. Мерфі.

Прогнозування динаміки часових рядів складних систем досліджувались вітчизняними та зарубіжними вченими. Відмітимо праці Н. Вінера, В.В. Леонтьєва, О.М. Дж. Кендала, а також українських вчених П. І. Бідюка, І. В. Баклана, Ю.П. Зайченка, М.З. Згуровського, О.Г. Івахненка, Н.К. Максишко, Н.Д. Панкратової, В.С. Степашка, О.І. Черняка та багатьох інших.

При великій кількості важливих робіт, широті досліджень, різноманітності одержаних в прогнозуванні результатів, все ще існують розділи прогностичної науки, в яких нові методи можуть покращити рішення, зробити його більш універсальним, конструктивним і точним. Необхідно зазначити, що останні десятиріччя - це початок активного

вивчення та переосмислення парадигм математичного моделювання економічних процесів. Переглядаються закони лінійної парадигми, ряд досліджень показують, що часові ряди складних економічних процесів не відповідають нормальному закону розподілу. Це, в свою чергу, ставить питання про неправомірність застосування широковідомих класичних методів прогнозування еволюційних процесів. В контексті економічних теорій з'являється економічна синергетика, як наука, один із підрозділів якої є теорія хаосу в поведінці економічних процесів. Дослідженню цих питань присвячені роботи як зарубіжних, так і вітчизняних авторів: В.-Б. Занг, Б. Мандельброт, Е.Петерс, І. Пригожин, С.П.Курдюмов, Г.Г. Малінецький.

Відносно молодим напрямом є еконофізика (Р. Мантен'я, Є. Стенлі). Вчені даного напрямку пропонують використовувати принципи, методи та моделі фізики для дослідження економічних процесів. Еконофізичний напрям активно розвивається українськими вченими (Я. І. Виклюк, В. Д. Дербенцев, В. М. Сапцін, О. А. Сердюк, В. М. Соловйов) та іншими. Роботи вищезазначених учених тісно зв'язані з задачами прогнозування економічних систем, зокрема динаміки фінансових ринків.

Вже завоювали визнання нейромережеві технології. Практика використання нейромереж показала їх ефективність в таких областях, як прогнозування, виявлення залежностей, класифікація, побудова систем підтримки прийняття рішень та інші. Вищезазначені технології також можуть застосовуватися на фінансових ринках, що дозволяє виявити та одержати нові знання про закономірності в динаміці індексів цінних паперів, про зміни показників економічної активності та про коливання обмінного курсу валют, котирувань акцій, облігацій та інших цінних паперів.

Економетричні методи прогнозування базуються на апроксимації часового ряду аналітичною функцією методом найменших квадратів.

Апроксимуюча функція не має властивості так званої “довгої пам’яті”, тому метод не є достатньо ефективним для прогнозування рядів зі змінною волатильністю. Методики, основані на технології нейронних мереж, показали високу ефективність у прогнозуванні часових рядів сучасних фінансових ринків, але вимагають спеціально підготовленої навчальної вибірки даних, яка повинна бути досить великою. Підготувати таку вибірку для фінансових рядів може бути проблематичним через брак даних високої часової дискретизації.

Запропонований метод прогнозування часових рядів на основі складних ланцюгів Маркова поєднує принципи практично всіх вищезгаданих методів, включаючи властивості довгої та короткої пам’яті в регресійні моделі, додаючи елементи фрактальності у неперервні економетричні моделі, а також суттєво доповнюючи сучасний нейромережевий підхід. За допомогою ієрархії часів методика дозволяє прогнозувати часові ряди, максимально використовуючи інформацію, представлену в ряді, дозволяє врахувати закономірності на усіх частотних рівнях.

Зв’язок роботи з науковими програмами, планами, темами.

Робота виконана у рамках дослідницької теми “Моделювання складних фінансово-економічних систем”, що проводяться на кафедрі економічної кібернетики Черкаського національного університету імені Богдана Хмельницького (номер державної реєстрації 0114U001026) у якості виконавця теми.

Мета і задачі дослідження. Метою даної роботи є розробка моделей часових рядів на основі складних ланцюгів Маркова для покращення точності довгострокових прогнозів часових рядів технічних, фінансово-економічних, соціальних та інших складних систем.

Для досягнення мети поставлено наступні задачі:

- провести аналіз сучасних методів прогнозування, здійснити їх

класифікацію для дослідження видів закономірностей в наборах даних, які в них використовуються;

- висвітлити проблеми сучасних трендових та лагових методів прогнозування та можливі шляхи їх подолання;
- розробити модель часового ряду на основі складних ланцюгів Маркова та на її основі розробити метод прогнозування;
- розробити методи оцінки можливих сценаріїв ряду для урахування невизначеності моделі;
- здійснити апробацію методу на часових рядах різної природи, оцінити можливий горизонт прогнозу та довірчі інтервали;
- виконати експериментальне порівняння якості прогнозів з відомими методами прогнозування індексів фондових ринків.

Об'єкт дослідження - процес прогнозування часових рядів при аналізі динаміки складних технічних, соціальних та фінансово-економічних систем.

Предмет дослідження - методи та моделі прогнозування часових рядів, оснований на складних ланцюгах Маркова.

Методи дослідження. У роботі використовувалися методи теорії випадкових процесів, складних ланцюгів Маркова, що дозволяють побудувати модель ряду; статистики та економетрії, які використані для вибору множини лагових змінних, які впливають на прогнозований показник, та для оцінювання матриць ймовірностей переходів між станами; методів нелінійної оптимізації для побудови нелінійної трендової моделі з гармонічними складовими для прогнозування низькочастотної складової часового ряду.

Інформаційну базу дослідження складають статистичні дані індексів світових фондових ринків, зведені з денною дискретизацією, доступні для вільного завантаження з мережі інтернет [1, 2] за приблизно 50 останніх

років, а також дані динаміки об'ємів газоспоживання з щохвилинною дискретизацією за 5 років.

Наукова новизна одержаних результатів. Основними науковими результатами, що виносяться на захист, є такі: Вперше:

- запропоновано метод прогнозування часових рядів на основі складних ланцюгів Маркова, тобто ланцюгів Маркова з пам'яттю, які, на відміну від відомих моделей часових рядів, використовують дискретне представлення часового ряду та враховують ефект пам'яті при такому представленні даних та здатні виявити закономірності ряду, приховані від класичних методів прогнозування;
- – вперше запропоновано метод прогнозування згідно ієрархії приростів часу та процедуру “склеювання”, що дозволило відтворювати мультифрактальні властивості часових рядів складних фінансово-економічних систем;

Удосконалено:

- методи дискретизації часового ряду, враховуючи особливості негаусового розподілу прибутковостей фінансових часових рядів;
- здійснено оптимізацію швидкодії алгоритму за рахунок використання хеш-таблиць;
- методи апроксимації та екстраполяції часових рядів на основі виокремлення низькочастотного тренду за допомогою дискретного Фур'є-продовження з двохчастотним наближенням;

Дістали подальшого розвитку:

- клітинно-автоматний підхід до прогнозування часових рядів введенням методу прогнозування при змінній довжині пам'яті.
- мультиагентний підхід до паралельного обчислення прогнозів часових рядів.

Практичне значення одержаних результатів. Запропоновані моделі часових рядів та алгоритми прогнозування можуть бути використані при аналізі та прогнозуванні часових рядів різної природи, зокрема діяльності фінансово-економічних організацій, яким необхідне прогнозування. Для технічних систем, що характеризуються динамікою показників у вигляді часових рядів, розроблені методи також можуть бути корисними для аналізу та прогнозування.

Особистий внесок здобувача. Основні результати, що виносяться на захист, отримані автором самостійно. Зі статті [3], опублікованої у співавторстві, автору належить розробка алгоритму та програмного забезпечення для прогнозування часових рядів.

Апробація результатів дисертації. Основні результати дослідження доповідалися на наукових конференціях різного рівня та наукових семінарах. Це такі конференції:

- Working Group on Physics of Socio-economic Systems, Dresden, October 5-8, 2009.;
- International Conference “Information Technologies, Management and Society”. April 16-17, Riga, Latvia.;
- X Міжнародної науково-технічна конференція “Системний аналіз та інформаційні технології”, м Київ, 20–24 травня 2008 р.; 21–26 травня 2009 р.; 21–26 травня 2010 р.; травень 2012 р. ;
- Проблеми інформатики та моделювання, м. Харків, 15-17 вересня 2010 р.
- VI Міжнародної науково-технічної конференції “Комп’ютерні технології в будівництві”, Київ-Севастополь, 9-12 вересня 2008 р.
- Гумбольдт-коллеґ 2009.
- Міжнародній науково-практичній конференції “Dynamical systems modeling and stability investigation”, DSMSI-2011, м. Київ.

- I Міжнародній науково-практичній конференції "Обчислювальний інтелект-2011", м. Черкаси,
- Міжнародній науково-практичній конференції "Ітонт-2009", "Ітонт-2012", м. Черкаси.
- Моніторинг, моделювання та менеджмент емерджентної економіки, 2008, м. Черкаси, 2010, м. Одеса.
- Інформаційні технології та моделювання в економіці. 2009 р., м. Черкаси.

Результати дисертаційної роботи були обговорені на наступних семінарах:

- Міжміський постійно діючий семінар "Моделювання складних систем", кафедра економічної кібернетики Черкаського національного університету імені Богдана Хмельницького.
- Семінар на кафедрі комп'ютерних технологій ЧДТУ
- Семінар на кафедрі інформатики і прикладної математики КДПУ

Публікації. Основні результати роботи викладено у 4-х статтях [4–6, 15], опублікованих у виданнях, що внесені до переліку наукових фахових видань України, 1-й статті [3] у зарубіжному фаховому журналі та додатково відображено в матеріалах конференцій [7–14, 16, 17].

Зміст роботи. Дисертація складається зі вступу, 4-х розділів, списку літератури (112 джерел) та додатків. Загальний об'єм дисертації складає 156 сторінок.

РОЗДІЛ 1

АНАЛІЗ ПРОБЛЕМИ ПРОГНОЗУВАННЯ ЧАСОВИХ РЯДІВ

У даному розділі проведено аналіз загальновідомих методів прогнозування часових рядів. Досліджено сучасні підходи до аналізу та прогнозування часових рядів, такі як нейронні мережі, нечіткі множини, нелінійна динаміка та ін.

1.1. Огляд та класифікація методів прогнозування часових рядів

Існує два головних напрями аналізу фінансових ринків: фундаментальний та технічний аналіз. Фундаментальний аналіз полягає у вивченні факторів, які впливають на стан економіки, секторів промисловості та певних компаній. Наприклад, для передбачення майбутніх цін на окремий вид акцій фундаментальний аналітик поєднує дослідження стану економіки, секторів промисловості та окремих компаній для визначення “справедливої” ціни (фундаментального значення), яка відображає реальну вартість підприємства та надає можливість передбачення майбутніх змін цієї вартості. Якщо фактична ціна акцій відхиляється від фундаментального значення, можна стверджувати, що акції або недооцінені, або переоцінені, а ринкова ціна коливається навколо фундаментального значення.

Існують деякі недоліки, притаманні фундаментальному аналізу: потреба значного часу для збору даних та розрахунків, що зумовлює відставання прогнозу від потоку інформації, яка постійно з’являється, а також завжди має місце суб’єктивність думок аналітиків.

Технічний аналіз базується на дослідженні динаміки змін ціни в минулому для передбачення її майбутнього значення. У цьому дослідженні цей підхід є основним, а прогнози базуються на динаміці денних цін акцій.

Теоретичні основи технічного аналізу базуються на твердженні, що зміни ціни відображають усю важливу інформацію, яка доступна на ринку. Цей факт впливає з припущення неефективності фінансових ринків. Гіпотеза ефективного ринку розглянута в роботах [18, 19]. Невдовзі після появи, гіпотеза ефективного ринку була піддана гострій критиці як теоретиками, так і емпіричними дослідниками.

Паралельно фундаментальному аналізу в даний час розвивається технічний аналіз – метод вивчення рухів цін, головним інструментом якого служать графіки [20–23]. Своє коріння сучасний технічний аналіз бере з теорії Чарльза Доу, впливаючи з неї прямо або побічно. Технічний аналіз увібрав у себе такі принципи і поняття, як напрям руху цін, (“ціни враховують усю відому інформацію”), збіжність і розбіжність ковзних середніх, динаміка об’ємів як дзеркало цінових змін, а також лінії підтримки та опору.

Розглянемо деякі підходи до класифікації сучасних методів прогнозування.

В роботах [24, 25] запропоновано класифікацію, в якій за джерелами прогнозованої інформації усі методи прогнозування поділяються на інтуїтивні та формалізовані. Для кожного методу з групи формалізованих будується математична модель, властивості якої використовуються при прогнозуванні. Пропонується поділ формалізованих методів на методи екстраполяції та методи моделювання систем [24, 25]. Інтуїтивні методи рекомендується використовувати для довгострокового прогнозування, коли закономірності розвитку важко формалізувати у вигляді математичної або комп’ютерної (алгоритмічної, імітаційної) моделі. Джерелом прогнозованої інформації для інтуїтивних методів

прогнозування є досвід експертів, тому інтуїтивні методи також варто називати експертними [25].

У роботах [26] розглядається структурний та формалізований підхід до побудови прогнозних моделей. Під структурним підходом розуміють дослідження закономірностей безпосередньо об'єкта, який продукує часові ряди, а під формалізованим – побудова моделей часових рядів, які є похідними певного процесу, суть якого невідома (концепція чорної скриньки).

Означення 1.1. *Часовим рядом* називається послідовність рівнів ряду

$$\{y_i\} = y_1, y_2, \dots, y_n \quad (1.1)$$

де i – натуральні числа, які відповідають моментам чи проміжкам вимірювання досліджуваної величини. Вимірювання відбуваються через рівні проміжки часу. y_i – рівні ряду.

Усі формалізовані методи прогнозування можуть бути подані у вигляді наступної математичної моделі:

$$y_{t+1} = f(t, y_t, \dots, y_{t-m}, x_t^{(1)}, \dots, x_{t-m_1}^{(1)}, x_t^{(2)}, \dots, x_{t-m_2}^{(2)}, x_t^{(k)}, \dots, x_{t-m_k}^{(k)}), \quad (1.2)$$

де y_t – величина показника, що прогнозується в момент часу t ; $x_t^{(i)}$ – величина i -го фактора в момент часу t , який впливає на y ; $m, m_1, \dots, m_k$ – довжини “пам’яті” ряду, які використовуються при прогнозуванні.

Виділимо, зокрема, недоліки загальноприйнятих класифікацій:

1. відсутність поділу моделей на лінійні та нелінійні;
2. ігнорування сучасних підходів до прогнозування часових рядів, таких як нейронні мережі, теорія нечітких множин, адаптивні фільтри тощо.

Нами пропонується удосконалена класифікація методів прогнозування, яка наведена на рисунку 1.1.

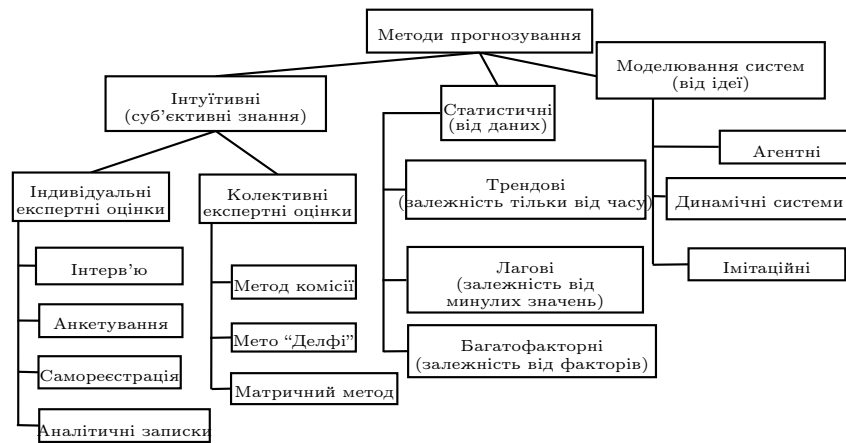


Рис. 1.1. Класифікація методів прогнозування

За основу класифікації взято об'єкт, вихідні дані чи модель якого дозволяє побудувати прогноз. У першій, експертній групі методів таким об'єктом є експерт, який опитується чи анкетується. Основна ідея даних методів полягає у проведенні опитування та статистичного опрацювання суб'єктивних прогнозів експертів. Оскільки формально джерелами знань є експерти, які володіють власними, здебільшого інтуїтивними правилами прийняття рішень, а експертні методи не будують ніяких моделей міркування експертів, а лише збирають дані, то експертні методи не є предметом дослідження даної книги. Але ідея дослідження колективної думки, оцінювання ваги, актуальності, "правильності", компетентності думок кожного експерта є дуже плідною. Майбутній розвиток цієї ідеї ми вбачаємо у побудові системи видобування знань із джерел новин мережі Інтернет, аналізу правдоподібності будь-якої інформації, включаючи гіпотези та чутки, саме засобами інформаційних технологій.

Одним із нових напрямів дослідження є теорія несилової взаємодії, запропонована у монографії [27]. Автором пропонується методи прогнозування, які базуються на цьому принципі існування природи та суспільства. У рамках даної роботи вказаний напрям розглядатись не буде через його складність.

Основна увага в даній роботі приділяється формалізованим методам прогнозування. Згідно з запропонованою нами класифікацією (рис. 1.1) дані методи поділяються на статистичні методи та методи моделювання систем. Методи з групи “Моделювання систем” направлені на побудову моделей, які описують безпосередньо функціонування досліджуваного об’єкту чи складної системи. Статистичні методи тут треба розуміти як методи, направлені на дослідження тільки похідних показників досліджуваного об’єкту, представлених у вигляді часових рядів.

Методи прогнозування з групи “Моделювання систем” базуються на моделях об’єкту дослідження. До них можна віднести імітаційне моделювання, під яким традиційно розуміється генерування послідовності випадкових подій чи реалізації випадкового процесу із заданими статистичними властивостями (розподілом ймовірності, автоковаріаційною функцією тощо) та оцінювання ризику за допомогою імітування критичних ситуацій у досліджуваній системі. Динамічні системи – це такі моделі, в яких закономірності розвитку об’єкта моделюються на основі диференціальних чи різницевих рівнянь. Якщо розглядати моделі з класу “динамічні системи” тільки як розв’язки диференціальних рівнянь (не звертаючи уваги на закономірності, які описують динамічні системи), то формально дані моделі можна віднести до статистичних, зокрема трендових. У цьому розумінні, розв’язок диференціального рівняння апроксимує досліджуваний ряд. Але в даному випадку акцент робиться на системні закономірності, на яких основана побудова диференціального рівняння, що описує динаміку системи. Тому динамічні системи, на нашу думку, доцільно вважати моделями самої системи.

Потужними можливостями для прогнозування складних систем, зокрема, фінансово-економічних, є агентні моделі [28–32]. Ідея їх побудови аналогічна моделям молекулярної динаміки у фізиці. Агентна модель є програмною системою, яка включає модель агентів, як окремих відносно

простих підсистем досліджуваної системи, та моделі середовища їх взаємодії. Прикладом може бути система підприємств певної галузі, розташованих на певній території, які певним чином взаємодіють між собою: множина трейдерів, які торгують на фінансовому ринку тощо. При взаємодії великої кількості агентів, з'являються емерджентні властивості побудованої системи, які безпосередньо не впливають із властивостей окремих агентів. Тому навіть прості моделі поведінки агента можуть пояснювати емерджентні явища в сучасних складних системах. На шляху ефективної побудови та використання агентних моделей існують певні проблеми: недостатня складність у поведінці агентів, неможливість точного копіювання усіх учасників системи, проблеми калібровки, налаштування невизначених параметрів алгоритму поведінки агента та інші. Тому це питання не входить до компетенції даної роботи.

Основна увага нашої роботи направлена на удосконалення статистичних методів прогнозування. Термін “статистичні” треба розуміти у широкому сенсі, як будь-які моделі одновимірних чи багатовимірних часових рядів. Традиційно в даній групі згадуються трендові, економетричні, авторегресійні моделі. Загальноприйнятою для статистичних методів моделлю, яка описує залежність, вважається лінійна, або ті, які зводяться до лінійної. Але дослідження показують, що як нелінійні аналітичні моделі, так і сучасні нейромережеві, нечіткі моделі показують значно кращі можливості у виявленні та оцінюванні складних залежностей. Тому всі методи групи “Статистичні” піддаються додатковій класифікації залежно від моделі залежності.

Відмінності різних методів будуть стосуватися структури вхідних даних x та виду функції $f(...)$ у відображенні 1.2. Таким чином, можна виділити такі види залежностей, які використовуються у сучасних методах прогнозування:

1. лінійна,

2. поліноміальна,
3. нейромережева,
4. нечітка,
5. випадкова.

В наступних розділах даної роботи будуть розглядатись вищезазначені залежності, що моделюються рядом методів прогнозування.

1.2. Перспективи використання трендових моделей при побудові сучасних методів прогнозування

Технічний аналіз [20–23] полягає у вивченні динаміки руху цін у минулому з метою визначення ймовірного напрямку їх розвитку в майбутньому. Поточна динаміка цін (тобто поточні очікування) порівнюється із зіставною динамікою цін у минулому, за допомогою чого досягається більш менш реалістичний прогноз.

Незалежно від того, якому типу аналізу віддає перевагу трейдер (особа, яка здійснює купівлю та продаж активів на фінансових ринках) [22] – фундаментальному або технічному, при складанні ринкових прогнозів він звертається до поняття “лінія тренду” (trendline). Поняття “лінія тренду” часто трактується неоднозначно і непослідовно. Враховуючи фрактальність, на різних часових рівнях може існувати декілька ліній тренду, які адекватно описують динаміку, що сформувалася. Тому при дослідженні трендів обов’язково необхідно включати поняття “горизонт прогнозування” та “інтервал трендостійкості” тощо.

Посібники по технічному аналізу пропонують методику вибору двох критичних точок, через які рекомендується проведення лінії тренду [22]. Графічний аналіз ліній тренду втратив минулу суб’єктивність і перетворився на суто механічну процедуру. З’явилися чіткі критерії проривів лінії тренду і можливість легко розрахувати цінові орієнтири, що

є достатнім для створення повноцінних торгівельних систем. На рисунках 1.2 і 1.3 представлено 2 лінії тренду: пропозиція спадаючою лінією, попит – зростаючою.

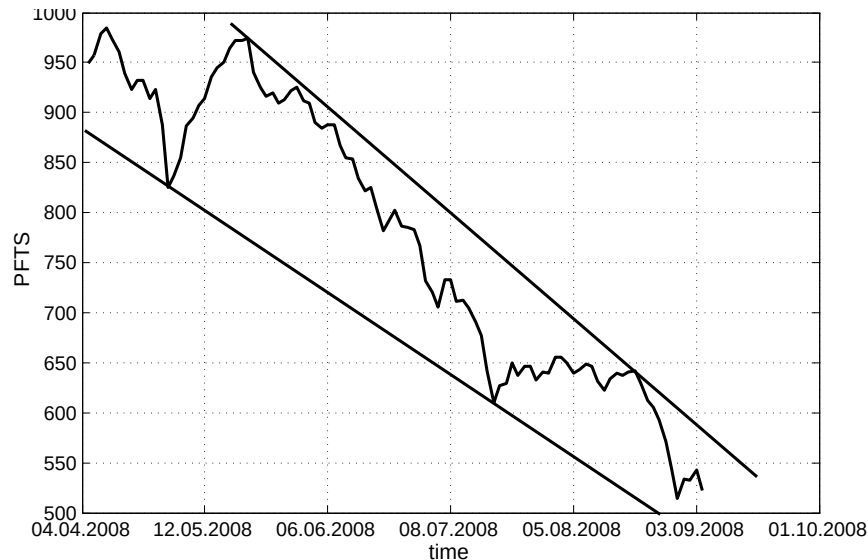


Рис. 1.2. Графік лінії тренду (низхідна лінія)

Можна звернути увагу, що на рисунку 1.2 поступове зниження цін відображене низхідною лінією “пропозиції”: цінові списи і западини також послідовно знижуються.

На рисунку 1.3 поступове підвищення цін відображене висхідною лінією “попиту”: цінові списи і западини також послідовно підвищуються. Складність полягає у виборі особливих точок, через які проходять ці прямі лінії (див. рис. 1.3). Як правило, аналітик привносить неабияку частку суб’єктивізму при побудові ліній тренду. Так, рух цін на ринку прийнято розглядати у ретроспективі – від минулого до майбутнього, тому і дати на графіку відкладаються зліва направо.

Загальноприйнята процедура побудови множини ліній тренду, з яких одна неухильно вважається вірною, також грішить неточністю та певною неувагою до дрібних деталей та фрактальних властивостей часового ряду.

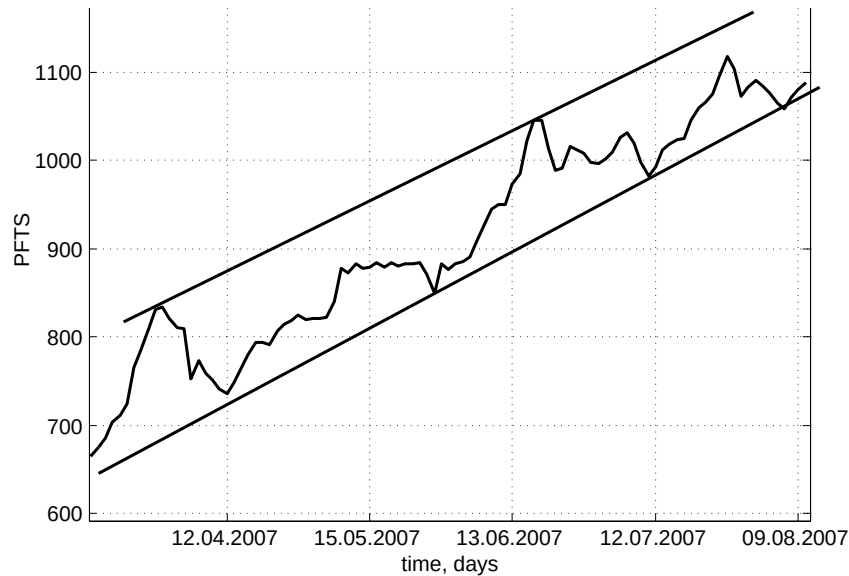


Рис. 1.3. Графік лінії тренду (висхідна лінія)

Тим часом успішне застосування ліній тренду вимагає від аналітика підвищеної уважності і послідовності.

Слід зазначити, що найбільш популярні класичні методи, які використовуються у світі для аналізу і прогнозування фінансового стану підприємств, пов'язані із застосуванням економетричних моделей, кожен змінну яких оцінює певний статистичний індикатор з тією чи іншою точністю, який вимірює певну сторону господарської діяльності досліджуваного об'єкту. У ряді випадків вони дозволяють виявити конкретні кількісні взаємозв'язки економічних процесів, що протікають на підприємствах-емітентах, з індикативними показниками, доступними для інвесторів. Проте в сучасних умовах використання розроблених на основі західної практики моделей за очевидних причин ускладнене.

Моделі часового ряду, в яких динаміка середнього значення ряду описується аналітичною залежністю тільки від часу, називаються трендовими моделями [25, 33, 34]. У більшості робіт [25, 33, 34] під трендовими моделями мають на увазі лінійні моделі тренду або ті, що до них зводяться.

Означення 1.2. *Трендовою моделлю* будемо називати аналітичну функцію виду

$$y = f(t, \vec{a}), \quad (1.3)$$

де t – час, заданий в певних одиницях (для даного дослідження – у днях); $\vec{a}$ – вектор параметрів, при яких модель найточніше описує залежність досліджуваного ряду $\{y_i\}$ від часу t та забезпечує мінімальні відхилення фактичних та змодельованих рівнів:

$$F(\vec{a}) = \|y_t - f(t, \vec{a})\| \quad (1.4)$$

на певному інтервалі часу $t : t_{first} < t < t_{last}$.

Зауважимо, що функція $f(t, \vec{a})$ може бути не тільки традиційною аналітичною (поліномом), але й гармонічною [13, 35, 36], складною нейронною мережею тощо.

Для побудови даних моделей здебільшого використовується метод найменших квадратів [25], або його модифікації для більш складних випадків.

До недоліків традиційних підходів варто віднести зведення до методу найменших квадратів, ігноруючи інші методи нелінійної чи стохастичної оптимізації, які у випадку багатьох екстремумів є більш ефективними. Також необхідно зазначити:

1) Апроксимація окремо параметрів тренду та гармонічних складових. Наші роботи показують, що оптимізація (мінімізація функціоналу нев'язки) для усіх параметрів разом, крім того декількох гармонік одночасно, дозволяють дещо покращити результати як апроксимації, так і прогнозу. Важливим розвитком таких моделей є методи, які враховують зміну частот;

2) Динамічні системи як трендові моделі. Розв'язок диференціальних рівнянь, заданих моделлю динамічної системи у вигляді залежності $y =$

$f(t)$ може вважатись трендовою у випадку, коли цей розв'язок є функцією тільки від одного параметру – часу.

Для сучасних методів прогнозування трендові моделі відіграють роль функції для нормалізації ряду, що дозволяє збільшити повторюваність навчальної вибірки.

Розглянемо застосування методу найменших квадратів, який часто використовується для оцінювання параметрів трендових моделей та для обчислення довірчих інтервалів прогнозних значень. Строго обґрунтування методу дано Марковим та Колмогоровим.

Для одновимірної випадку

$$y'(t) = \alpha + \beta \cdot x_t + u_t; \quad L = \sum_{t=1}^T (y'_t - y_t)^2.$$

Для багатовимірної випадку:

$$y = a_0 + a_1 \cdot x_1 + a_2 \cdot x_2 + a_3 \cdot x_3 + \dots + a_k \cdot x_k + u_k; \quad Y = X \cdot A + U;$$

Вектор невідомих параметрів $A = (a_0, a_1, a_2, \dots, a_k)$ оцінюється за формулою:

$$A = (X^T \cdot X)^{-1} \cdot X^T \cdot Y.$$

В роботах [24, 25, 34] даний метод використовується для побудови багатофакторних та авторегресійних прогнозних моделей.

Передумови використання даного методу базуються на твердженнях

- незміщеності оцінок параметрів;
- відсутності мультиколінеарності векторів вхідних параметрів;
- відсутності автокорельованості залишків;

Для подолання таких труднощів рекомендуються процедури нормалізації і стандартизації даних. Метод найменших квадратів може бути також застосований для деяких видів нелінійних трендових моделей,

які описані у роботі [25]. Їх можна звести до лінійних за допомогою заміни змінних.

Варто зауважити, що оцінювання коефіцієнту детермінації та статистичної значущості результатів здійснюється не для вихідних змінних, а для лінеаризованих, що обов'язково необхідно враховувати.

Розглянемо обчислення довірчих інтервалів для прогнозу згідно трендової моделі. Для нормального розподілу $N(\mu, \sigma)$ з невідомими параметрами μ та σ процес відповідає наступній допоміжній статистиці [37] :

$$\frac{X_{n+1} - \bar{X}_n}{S_n \sqrt{1 + 1/n}} \sim T^{n-1}.$$

Оцінений розподіл для майбутнього значення X_{n+1} є розподілом прогнозного значення

$$\bar{X}_n + S_n \sqrt{1 + 1/n} \cdot T^{n-1}.$$

Ймовірність того, що значення X_{n+1} буде знаходитись у заданому інтервалі, наступна:

$$\Pr \left(\bar{X}_n - T_a S_n \sqrt{1 + (1/n)} \leq X_{n+1} \leq \bar{X}_n + T_a S_n \sqrt{1 + (1/n)} \right) = p,$$

де $T_a - 100((1+p)/2)$ -на квантиль t -розподілу Стюдента з $(n-1)$ ступенями свободи. Тому величини

$$\bar{X}_n \pm T_a S_n \sqrt{1 + (1/n)}$$

є граничними точками 100ρ %-го прогнозного інтервалу для значення X_{n+1} .

На нашу думку, до трендових моделей необхідно віднести також лінійно-гармонічну модель, наведену у монографії [35]. Розглянемо її.

$$y(t) = c_0 + c_1 t + \sum_{i=1}^n \left(a_i \cos \frac{2\pi}{T_i} t + b_i \sin \frac{2\pi}{T_i} t \right), \quad (1.5)$$

де $t = 1, 2, \dots, n$ – час, c_0, c_1, a_i, b_i, T_i – параметри моделі. Для оцінювання параметрів пропонується сканування діапазону можливих періодів T_i та оцінювання значень c_0, c_1, a_i, b_i параметрів багатofакторної моделі з факторами $(t, \cos \frac{2\pi}{T_i} t, \sin \frac{2\pi}{T_i} t)$ методом найменших квадратів. З множини перебраних періодів T_i вибирається той, який забезпечує мінімальне значення функціоналу нев'язки

$$Z(c_0, c_1, a_i, b_i) = \sum_{t=T_{min}}^{T_{max}} \left[Y(t) - c_0 - c_1 t - a_i \cos \frac{2\pi}{T_i} t + b_i \sin \frac{2\pi}{T_i} t \right]^2 \rightarrow \min. \quad (1.6)$$

Даний метод показує достатньо високу точність прогнозів для рядів урожайності зернових культур України. До недоліків даної моделі можна віднести наступні:

- повний перебір частот не є оптимальним, такий підхід не дає можливості точно оцінити період коливань;
- неможливість оптимізації моделі за параметром T_i ;
- можливі проблеми з визначенням параметрів за МНК (мультиколінеарність, гетероскедастичність тощо).

Удосконалення даного методу розглядається в розділі 2.5 даної роботи.

Відомим методом аналізу частотного спектру часового ряду є метод “Гусениця”, більш відомий в західних публікаціях як метод сингулярного спектрального аналізу (Singular Spectrum Analysis, SSA) [38–40]. Суть його полягає у конструюванні траєкторної матриці з лагових змінних часового ряду та аналізу головних компонент отриманої траєкторної матриці. Розглянемо суть даного методу.

Нехай задано часовий ряд $\{y_t\}$. Метод “Гусениця” складається з наступних кроків.

1. Нормалізація вхідного ряду

$$y_t^* = y_t - \bar{y}_t,$$

або

$$y_t^* = \frac{y_t - \bar{y}_t}{\sigma},$$

де $\bar{y}_t$ - середнє значення ряду, σ - середньоквадратичне відхилення. 2.

Побудова траєкторної матриці

$$X = \begin{pmatrix} y_1^* & y_2^* & \cdots & y_m^* \\ y_2^* & y_3^* & \cdots & y_{m+1}^* \\ \cdots & & & \cdots \\ y_{N-m}^* & y_{N-m+1}^* & \cdots & y_N^* \end{pmatrix},$$

де m - розмірність вектору вкладення. Матриця X називається ганкелевою.

3. Побудова кореляційної матриці

$$C = \frac{1}{m} X^T X.$$

4. Обчислення власних векторів та власних значень матриці C . Нехай $v^{(1)}, v^{(2)}, \dots, v^{(m)}$ - власні вектори матриці C , а $\lambda_1, \lambda_2, \dots, \lambda_m$ - її власні значення.

5. Вибір мінімального базису власних векторів. З множини власних векторів пропонується вибрати підмножину векторів, які описують тільки суттєву варіацію часового ряду. Для цього упорядкуємо пари $(\lambda_i, v^{(i)})$ у порядку спадання власних значень. Нехай отримані номери власних векторів є наступними: $k_1, k_2, \dots, k_m$. Способи вибору підмножини власних векторів для аналізу часових рядів детально обговорюються в [41]. Розглянемо найпростіший метод. Визначимо τ як кількість найбільших власних векторів, питома вага власних значень яких більша, наприклад, за 75% ваги усіх власних значень.

Залишемо матрицю V_x у вигляді отриманих власних векторів-стовбців:

$$V_x = \left(v^{(k_1)}, v^{(k_2)}, \dots, v^{(k_\tau)} \right) = \begin{pmatrix} v_1^{(k_1)} & v_1^{(k_2)} & \dots & v_1^{(k_\tau)} \\ v_2^{(k_1)} & v_2^{(k_2)} & \dots & v_2^{(k_\tau)} \\ \dots & \dots & \dots & \dots \\ v_m^{(k_1)} & v_m^{(k_2)} & \dots & v_m^{(k_\tau)} \end{pmatrix}.$$

6. Відображення матриці X в новому базисі власних векторів:

$$U_x = V_x^T X = (U_1, \dots, U_\tau).$$

7. Відновлення вихідної матриці за першими r головними компонентами:

$$\tilde{X} = \left(v^{(k_1)}, v^{(k_2)}, \dots, v^{(k_r)} \right) \begin{pmatrix} U_1^T \\ \dots \\ U_r^T \end{pmatrix}.$$

8. Діагональне усереднення (ганкелізація) матриці $\tilde{X}$: Оскільки відновлена матриця не буде ганкелевою (елементи діагоналей, паралельні побічній діагоналі, не рівні між собою), необхідно здійснити зведення матриці $\tilde{X}$ до такого виду. Елементи нової матриці матимуть значення:

$$f_s = \begin{cases} \frac{1}{s} \sum_{i=1}^s \tilde{x}_{i,s-i+1}, 1 \leq s \leq \tau, \\ \frac{1}{\tau} \sum_{i=1}^{\tau} \tilde{x}_{i,s-i+1}, \tau \leq s \leq n, \\ \frac{1}{N-s+1} \sum_{i=1}^{N-s+1} \tilde{x}_{i+s-n,n-i+1}, \tau \leq s \leq N. \end{cases}$$

Оригінальний метод “Гусениця” був розроблений для виявлення та аналізу прихованих періодичностей в часових рядах, але не містив моделі часового ряду, яку можна використати для побудови прогнозів. У роботі [41] пропонується декілька підходів побудови прогнозних моделей на основі вище запропонованого сингулярного розкладу. Не всі вони можуть бути застосовані до реальних часових рядів, так як базуються на понятті строго

детермінованої послідовності. Розглянемо підхід, який найбільш ефективно може використовуватись для аналізу часових рядів зі стохастичними складовими.

Класична економетрія включає багато робіт з прогнозування часових рядів [34, 42–44]. Головна ідея економетричних методів полягає у наближенні часового ряду до лінійного виду за допомогою лінеаризації. Далі будується лінійна регресійна модель та використовується для передбачення майбутніх значень, $y_t = a + b \cdot y_{t-1} + u$, де a та b – регресійні параметри, u – залишки. Моделюючи відхилення на основі t -статистики, можна оцінити довірчі інтервали отриманого точкового прогнозу. Даний метод є більш грубий, але працює для детермінованих лінійних трендів.

1.3. Авторегресійні методи прогнозування

Означення 1.3. Процесом з авторегресією $AR(p)$ називається послідовність рівнів ряду $\{y_i\}$, яка задовольняє рівняння:

$$y_{n+1} = a_0 + a_1 y_n + a_2 y_{n-1} + \dots + a_{k-1} y_{n-k+1} + a_k y_{n-k} + e, \quad (1.7)$$

де $\{a_i\}$ – параметри авторегресії, k – порядок авторегресійної моделі, y_i – рівні ряду.

У монографіях [26, 45–47] пропонується метод знаходження аналітичного розв’язку авторегресійного рівняння (1.7) відносно параметру часу t , що дозволяє звести авторегресійну модель часового ряду до трендової. Дане перетворення базується на пошуку частинних розв’язків авторегресійного рівняння виду $y_i^h(k) = A_i \cdot (\alpha_i)^k$.

Для рівнянь порядку вище 2-го можна виконати аналогічні перетворення, але виникають труднощі, пов’язані з розв’язанням алгебраїчного рівняння високого степеня. На нашу думку, даний метод не покращує прогнозні можливості авторегресійних моделей у порівнянні

з ітераційним обчисленням прогнозованих значень ряду, але розв'язок авторегресійного рівняння відносно невідомої функції $y(t)$ дозволяє провести класифікацію можливих режимів авторегресійних процесів [26, 48].

Також варто зауважити, що авторегресійні випадкові процеси у ряді робіт [25, 49] називаються марківськими, так як описують процеси з короткою пам'ятю, тобто лінійну залежність майбутньої величини від ряду попередніх. Оскільки під дискретними ланцюгами Маркова розуміють випадкові процеси з пам'яттю один крок, то марківськими процесами у роботах [25, 49] називаються авторегресійні моделі порядку 1.

На нашу думку, клас марківських процесів не обмежується лінійними авторегресійними моделями, так як залежність може бути нелінійною та записуватись у вигляді алгоритму абстрактної машини або клітинного автомату [50].

Означення 1.4. *Процесом ковзного середнього $MA(q)$ називається послідовність рівнів ряду $\{y_i\}$, яка задовольняє рівняння:*

$$y_{n+1} = b_0 + b_1 e_n + b_2 e_{n-1} + \dots + b_{k-1} e_{k-2} + b_k e_{k-1}, \quad (1.8)$$

де $\{b_i\}$ – параметри авторегресії, k – порядок авторегресійної моделі, e_i – залишки.

Для моделювання та прогнозування частіше використовуються комбіновані $ARMA(p, q)$ -моделі.

Означення 1.5. *Процесом авторегресії – ковзного середнього $ARMA(p, q)$ називається послідовність рівнів ряду $\{y_i\}$, яка задовольняє рівняння:*

$$\begin{aligned} y_{n+1} = & a_0 + a_1 y_n + a_2 y_{n-1} + \dots + a_{k-1} y_{k-2} + a_k y_{k-1} + \\ & + b_0 + b_1 e_n + b_2 e_{n-1} + \dots + b_{k-1} e_{k-2} + b_k e_{k-1}, \end{aligned} \quad (1.9)$$

де $\{a_i\}, \{b_i\}$ – параметри авторегресії, p – порядок авторегресійної складової моделі, q – порядок складової ковзного середнього, e_i – залишки.

ARIMA-моделі охоплюють досить широкий спектр часових рядів, а незначні модифікації цих моделей дозволяють досить точно описувати часові ряди з сезонністю. Говорячи про ARIMA-моделі, матимемо на увазі, що окремими випадками ARIMA є AR-, MA- і ARMA-моделі. Крім того, виходитимемо з того, що вже здійснено підбір моделі для аналізованого часового ряду, включаючи ідентифікацію вибраної моделі.

Спрогнозуємо невідоме значення, x_{t+l} , $l > 1$ вважаючи, що x_t – останнє за часом спостереження аналізованого часового ряду, наявне в нашому розпорядженні і позначимо його через x_t^l . Відмітимо, що хоча x_t^l і x_{t-1}^{l+1} позначають прогноз одного і того ж невідомого значення x_{t+1} , але обчислюються вони по-різному, оскільки є вирішеннями різних завдань.

Ряд x_τ , що аналізується в рамках ARIMA (p, k, q) -моделі, представимо (при будь-якому $\tau > k$) у вигляді

$$\begin{aligned} (1 - \alpha_1 L - \dots - \alpha_p L^p) \cdot \sum_{j=0}^k (-1)^j C_k^j x_{\tau-1} = \\ = \delta_\tau - \theta_1 \delta_{\tau-1} - \dots - \theta_q \delta_{\tau-q}, \end{aligned} \quad (1.10)$$

де L – оператор зсуву функції часу на один часовий такт назад.

Зі співвідношення (1.10) можна виразити x_τ для будь-якого $\tau = t - q, \dots, t - 1, t + 1$. Отримуємо:

$$\begin{aligned} x_\tau = \left(\sum_{j=1}^p \alpha_j L^j \right) x_\tau + \left(1 - \sum_{j=1}^p \alpha_j L^j \right) \cdot \left(\sum_{i=0}^k (-1)^i C_k^i x_{\tau-i} \right) + \\ + \left(1 - \sum_{j=1}^q \theta_j L^j \right) \delta_\tau. \end{aligned} \quad (1.11)$$

Правими частинами цих співвідношень є лінійні комбінації з $p + k$ попередніх (по відношенню до лівої частини) значень аналізованого процесу x_τ , доповнені лінійними комбінаціями поточного і q попередніх значень випадкових залишків δ . Причому коефіцієнти, за допомогою яких

ці лінійні комбінації підраховуються, відомі, оскільки виражаються в термінах уже оцінених параметрів моделі. Цей факт і дає можливість використовувати співвідношення (1.11) для побудови прогнозних значень часового ряду на l тактів часу вперед. Теоретичну базу такого підходу до прогнозування забезпечує відомий результат, відповідно до якого найкращим (за середньоквадратичною похибкою) лінійним прогнозом у момент часу t з випередженням l є умовне математичне сподівання випадкової величини x_{t+l} , обчислене за умови, що всі значення x_τ до моменту часу t . Цей результат є окремим випадком загальної теорії прогнозування [51–53].

Таким чином, визначається наступна процедура побудови прогнозу по відомій до моменту прогнозування траєкторії часового ряду:

- 1) за формулою (1.11) обчислюються ретроспективні прогнози $x_{t-l}^l \dots x_{t-1}^l$ на базі попередніх значень часового ряду;
- 2) використовуючи формули (1.11) для $\tau > t$, підраховуються умовні математичні сподівання для обчислення прогнозних значень.

Метод експоненціального згладжування

Досить ефективним і надійним методом прогнозування є метод експоненціального згладжування [25, 53]. Основні переваги методу полягають у можливості оцінювання вагів початкової інформації, у простоті обчислювальних операцій, в гнучкості опису різних динамічних процесів. Метод експоненціального згладжування дає можливість отримати оцінку параметрів тренду, що характеризує не середній рівень процесу, а тенденцію, що склалася до моменту останнього спостереження. Після появи робіт Р. Брауна найчастіше даний метод застосовується для реалізації короткострокових і середньострокових прогнозів. Для методу експоненціального згладжування основним і найбільш важким моментом є вибір параметра згладжування α , початкових умов і степеня прогнозуючого полінома [25, 51].

Нехай початковий динамічний ряд описується функцією

$$y_t = a_0 + a_1 t + \frac{a_2}{2} t^2 + \dots + \frac{a_p}{p!} t^p + \varepsilon_t. \quad (1.12)$$

Метод експоненціального згладжування, що є узагальненням методу ковзного середнього, дозволяє побудувати такий опис процесу (1.12), при якому новим спостереженням додаються більші ваги в порівнянні зі старими, причому ваги спостережень спадають за експонентою. Вираз

$$S_t^{[k]}(y) = \alpha \sum_{i=0}^n (1 - \alpha)^i S_{t-1}^{[k-1]}(y)$$

називається експоненціальною середньою до n -го порядку для ряду y_t , де α – параметр згладжування.

У розрахунках для визначення експоненціальної середньої користуються рекурентною формулою [54]

$$S_t^{[k]}(y) = \alpha S_t^{[k-1]}(y) + (1 - \alpha) S_{t-1}^{[k]}(y)$$

Як було відмічено, важливу роль в методі експоненціального згладжування грає вибір оптимального параметра згладжування α , оскільки саме він визначає оцінки коефіцієнтів моделі, а отже, і результати прогнозу. Вибір параметра α доцільно пов'язувати з точністю прогнозу, тому для більш обгрунтованого вибору α можна використовувати процедуру узагальненого згладжування, яка дозволяє отримати співвідношення, що пов'язують дисперсію прогнозу і параметр згладжування [53]. Істотним для практичного використання є питання про вибір порядку прогнозуючого полінома, що багато в чому визначає якість прогнозу. У роботі [53] показано, що перевищення другого порядку моделі не приводить до істотного збільшення точності прогнозу, але значно ускладнює процедуру розрахунку.

Відзначимо, що даний метод є одним з найбільш популярних і широко застосовується в практиці прогнозування. Враховуючи, що метод експоненціального згладжування є узагальненням методу найменших квадратів, можна сподіватися, що він удосконалюватиметься і далі як в теоретичному, так і в прикладному аспекті. Один з найбільш перспективних напрямів розвитку даного методу є метод різницевого прогнозування, в якому саме експоненціальне згладжування розглядається як окремий випадок [53].

Огляд літератури дозволив виявити ряд недоліків загальновідомих методів прогнозування і завдань, які необхідно вирішити в даній області.

Головний недолік методу експоненціального згладжування полягає в тому, що часовий ряд розглядається ізольовано від інших явищ. Крім того, точність прогнозу помітно падає при довгостроковому прогнозуванні [53].

Слабким місцем усіх адаптивних методів, і методів експоненціального згладжування зокрема, є підбір відповідного до даного конкретного завдання параметра згладжування (адаптації) α . Навіть при оптимальному підборі параметра модель Брауна поступається в точності прогнозу ARIMA (0,1,1) -моделі.

Недоліками методу найменших квадратів є використання процедури оцінки, що передбачає обов'язкове виконання цілого ряду передумов, порушення яких може привести до значних помилок:

- 1) випадкові похибки повинні мати нульову середню та скінченні дисперсії і коваріації;
- 2) кожен вимір випадкової похибки характеризується нульовим середнім, не залежним від значень спостережуваних змінних;
- 3) дисперсії кожної випадкової похибки є постійними, їх величини незалежні від значень спостережуваних змінних (гомоскедастичність);
- 4) відсутність автокореляції похибок, тобто значення похибок різних спостережень мають бути незалежні один від одного;

5) випадкові похибки повинні відповідати нормальному закону розподілу;

6) значення ендогенної змінної x не містять похибок вимірювання і мають скінченні середні значення та дисперсії.

Вибір моделі у кожному конкретному випадку здійснюється за цілим рядом статистичних критеріїв, наприклад за дисперсією, кореляційним відношенням та ін. Класичний метод найменших квадратів передбачає рівноцінність початкової інформації в моделі. У реальній практиці майбутня поведінка процесу значно більшою мірою визначається більш новими спостереженнями, ніж минулими. Ця властивість вказує на зменшення цінності ранньої (більш віддаленої у часі, старої) інформації.

Важливим моментом отримання прогнозу за допомогою методу найменших квадратів є оцінка адекватності отриманого результату. Для цієї мети використовується цілий ряд статистичних характеристик:

- оцінка стандартної помилки;
- середня відносна помилка оцінки;
- середнє лінійне відхилення;
- кореляційне відношення для оцінки надійності моделі;
- оцінка достовірності вибраної моделі через значущість індексу кореляції за Z -критерієм Фішера;
- оцінка достовірності моделі за F -критерієм Фішера;
- перевірка на наявність автокореляцій за критерієм Дарбіна-Уотсона.

Недоліки методу найменших квадратів обумовлені жорсткою фіксацією тренду і надійний прогноз при цьому можна отримати тільки на невеликий період упередження. Тут мова йде про обмежені можливості методів математичної статистики, теорії розпізнавання образів, теорії випадкових процесів, нечітких методів та ін., оскільки багато реальних процесів не можуть адекватно бути описані за допомогою традиційних статистичних

моделей, оскільки є істотно нелінійними і мають або хаотичну, або квазіперіодичну, або змішану основу.

Недоліком методу адаптивного згладжування є те, що тільки при дуже довгих рядах можна отримати надійний прогноз на інтервал більший, ніж при звичайному експоненціальному згладжуванні. На жаль, для даного методу немає строгої процедури оцінки необхідної або достатньої довжини початкової інформації, для скінчених рядів немає конкретних умов оцінки точності прогнозу.

Проблеми з методом Бокса-Дженкінса (моделі авторегресії – ковзного середнього) пов'язані, перш за все, з неоднорідністю часових рядів і складністю практичної реалізації методу. В цілому, результати застосування традиційних технологій оцінювання і прогнозування фінансового стану компаній, а також реальної вартості пакетів їх цінних паперів (акції, облігації), які вільно продаються і купуються на фондовому ринку, можна назвати обмеженими. Обмеженість цих методів полягає у їх залежності від початкових умов і відсутності гнучкості. Жорсткі статистичні припущення про властивості часових рядів обмежують можливості методів математичної статистики, теорії розпізнавання образів, теорії випадкових процесів, нечітких методів та ін. Вони не здатні враховувати те, що відносна значущість окремих показників фінансової звітності та чинників, що їх визначають, на практиці змінюються з часом, іноді дуже різко і непередбачувано. Окрім цього, традиційні підходи характеризуються обмеженою інформативністю, оскільки призначені для опису якісних чинників або закономірностей в кількісних термінах. Таким чином, на зміну традиційним технологіям повинні прийти нові підходи, ефективніші в умовах структурної нестабільності світової економіки [55]. Останнім часом усе більше уваги приділяється дослідженню і прогнозуванню фінансових часових рядів з використанням теорії динамічних систем, теорії хаосу. Це досить нова область, яка

є популярним розділом математичних методів економіки, що активно розвивається [54, 56–59].

1.4. Використання нелінійних функцій, нейронних мереж та генетичних алгоритмів при побудові трендових та авторегресійних моделей.

Багато реальних процесів, у тому числі і показники ринку цінних паперів, не можуть адекватно бути описані за допомогою традиційних статистичних моделей, оскільки по суті є істотно нелінійними і мають або хаотичну, або квазіперіодичну, або змішану (стохастика + хаос-динаміка+детермінізм) основу [58–60]. У даній ситуації адекватним апаратом для вирішення завдань аналізу і прогнозування ринку цінних паперів можуть служити штучні нейронні мережі, що реалізуються на базі концепцій штучного інтелекту. Програмно-математичною основою цих методів є нейронні мережі, що самоорганізуються [61–66].

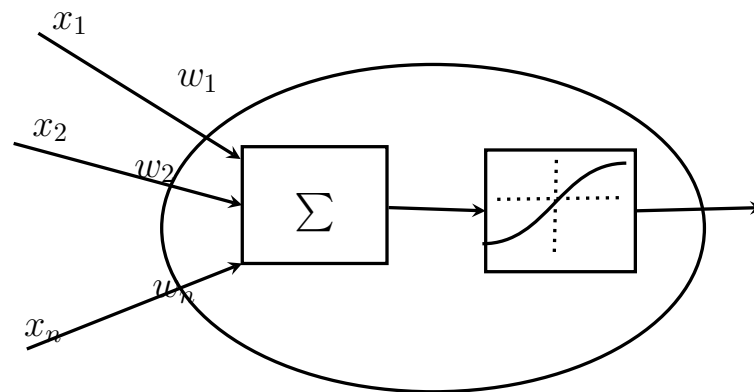


Рис. 1.4. Математична модель нейрона

Нейронна мережа – це сукупність нейронних елементів і зв’язків між ними (рис. 1.5). Основний елемент нейронної мережі – це формальний нейрон (рис. 1.4), що здійснює операцію нелінійного перетворення суми добутків вхідних сигналів на вагові коефіцієнти

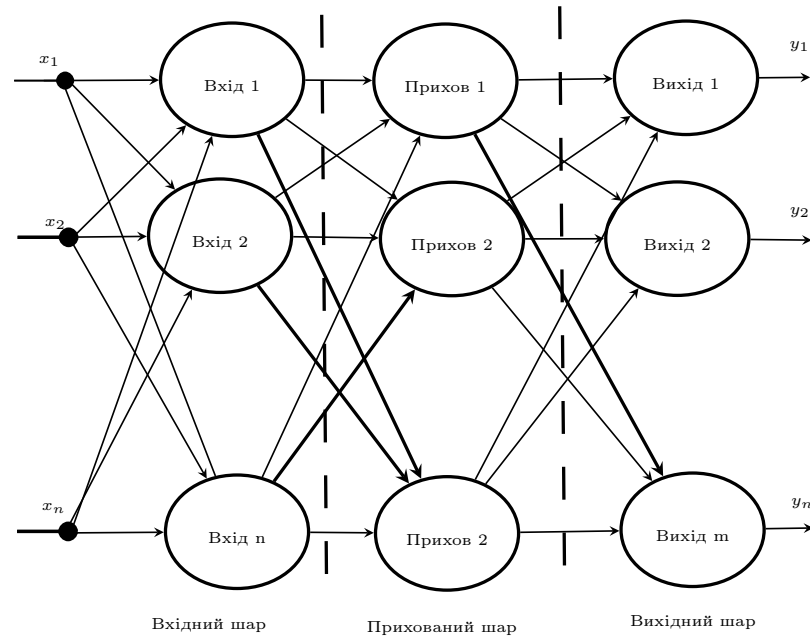


Рис. 1.5. Приклад нейронної мережі – багат шаровий персептрон

$$Y = F \left(\sum_{i=1}^n w_i \cdot x_i \right) \cdot F(WX), \quad (1.13)$$

де $X = (x_1, x_2, \dots, x_n)^T$ – вектор вхідного сигналу, $W = (w_1, w_2, \dots, w_n)$ – ваговою вектор; F – оператор нелінійного перетворення.

В даний час методи штучних нейронних мереж вже довели свою високу ефективність при дослідженні фінансово-економічних систем. Штучні нейронні мережі є апаратом з області нейрокомп'ютингу (neural computing), області обчислювальних технологій, що останнім часом швидко розвиваються. Обчислювальні операції в таких мережах виконуються великим числом порівняно простих процесорних елементів (processing element). Структура мережі (network) тотожна математично визначеній структурі обчислювальної системи (графу мережі), в якій усі операції виконуються в певних вузлах, а потік інформації відображається направленими ребрами графа. Кожен вузол (нейрон) мережі представлений процесорним елементом, нейронною клітиною, яка спільно з багатьма іншими процесорними елементами утворює

нейронну обчислювальну мережу. Аналогом такого вузла у фізіологічній нервовій системі є нервова клітина мозку [55]. У загальному випадку штучна нейронна мережа є адаптивною нелінійною динамічною системою. За допомогою рівноважних станів такої мережі можна вирішувати математичні або обчислювальні задачі, такі як апроксимація функцій, класифікація, кластеризація, прогнозування та багато інших.

Існує два класи нейронних мереж:

- мережі, що навчаються з учителем;
- мережі, що навчаються без учителя.

Нейронні мережі, що навчаються з учителем, є засобами для видобування з набору даних інформації про взаємозв'язки між входами і виходами нейромережі. Ці взаємозв'язки можуть бути записані математичними рівняннями, які можна використовувати для прогнозування або прийняття управлінських рішень. “Учителем” в даному випадку виступає набір параметрів, який дослідник розміщує на виході мережі. На вхід мережі при цьому подається відповідний даному виходу вхідний набір даних. Мережа навчається встановлювати залежність між початковою інформацією і результатами адаптивного ітераційного процесу [55].

Алгоритм роботи мережі, що навчається без учителя, ґрунтується на навчанні змаганнями. Забезпечується така поведінка нейронів мережі, що при подаванні чергового набору даних на вхід вони як би “змагаються” один з одним за якнайкращу відповідність вихідному набору згідно вибраним критеріям. В результаті змагання визначається нейрон переможець, після чого структура мережі піддається корекції. Клас нейронних мереж, що самоорганізуються, навчаються без вчителя, позначається терміном адаптивні нейронні мережі (АНМ). Важливою відмінністю методу АНМ є те, що він є чисельним, а не символьним методом обробки даних. Однією з унікальних особливостей АНМ є те, що вона надає внутрішньо

властиві їй точні і прості механізми для поділу обчислювального завдання на окремі потоки, що дозволяє проводити обчислення з високим ступенем паралельності. Навчання без учителя дає можливість виявляти у вхідних наборах даних невідомі раніше структури або закономірності, що відображає здатність АНМ до узагальнення (generalization) на основі вхідних прикладів. Ця властивість дозволяє узагальнювати великі набори багатовимірних даних, якими є фінансово-економічні показники підприємств, виявляти і демонструвати структури, що містяться в них, а також виявляти нові образи і взаємозв'язки в таких наборах даних. Нейронні мережі є потужним інструментом прогнозування. Закладені в них генетичні алгоритми, еволюціонуючи природним чином, дозволяють виявити оптимальну структуру нейронної мережі [67–69].

Перші кроки в області штучних нейронних мереж зробили в 1943 р. Уорен Мак-Калок та Уолтер Пітс. Вони показали, що за допомогою порогових нейронних елементів можна реалізувати числення будь-яких логічних функцій [64]. У 1949 р. Хебб запропонував правило навчання, яке стало математичною основою для навчання ряду нейронних мереж [63]. У 1957-1962 рр. Ф. Розенблатт запропонував і дослідив модель нейронної мережі, яку він назвав персептроном [65]. У 1959 р. В. Відроу і М. Хофф запропонували процедуру навчання для лінійного адаптивного елементу – AD ALINE. Процедура навчання отримала назву “Дельта правило” [64]. У 80-і роки значно розширюються дослідження в області нейронних мереж. Д. Хопфілд в 1982 р. дав аналіз стійкості нейронних мереж із зворотними зв'язками і запропонував використовувати їх для вирішення задач оптимізації. Т. Кохонен розробив і дослідив нейронні мережі, що самоорганізуються. Ряд авторів запропонували алгоритм зворотного поширення помилки, який став потужним засобом для навчання багат шарових нейронних мереж [63–65]. В даний час розроблено велике число нейросистем, які використовуються у різних

областях: прогнозуванні фінансових показників, управлінні, діагностиці в медицині та техніці, розпізнаванні образів тощо [70, 71].

Для навчання мережі використовуються різноманітні алгоритми навчання і їх модифікації [72–76].

Методам аналізу і підвищення якості навчальних вибірок для нейромережевих прогнозних моделей присвячена робота [77]. Запропоновані міри несуперечливості та повторюваності даних навчальної вибірки, а також методи підвищення якості навчальних вибірок, згідно запропонованим мірам. Методи покращення полягають у попередньому перетворенні даних (ППД) першого рівня, яке збільшує повторюваність даних (нормуванні, згладжуванні, заміні ряду абсолютними та відносними приростами, вибору множини вхідних сигналів), та ППД другого рівня, що забезпечує однозначність вибору розміру опису образу навчальних вимірів. Для цього здійснюється пошук розміру образу, при якому часовий ряд розбивається на ділянки різної довжини, складність на яких однакова. Для кожної ділянки формулюється стислий опис за допомогою спільної апроксимуючої функції. Порядок та клас апроксимуючої функції задає складність кожної ділянки. Пропонується модифікований метод формування навчальної вибірки, який, на відміну від методу “вікон”, дозволяє зняти неоднозначність вибору розміру розпізнаваного образу навчальних наборів.

Авторами [78, 79] розроблений алгоритм навчання мережі, який мінімізує середньоквадратичну помилку нейронної мережі за рахунок використання адаптивного кроку навчання $\tau(t)$. Пропонується використовувати логарифмічну функцію активації для вирішення завдань прогнозування, яка дозволяє отримати прогноз значно точніше. Аналіз багатоварових нейронних мереж авторів [78, 79] і алгоритмів їх навчання дозволив виявити ряд недоліків і проблем, що можуть виникнути:

- невизначеність у виборі числа шарів і кількості нейронних елементів у шарі;
- повільна збіжність градієнтного методу з постійним кроком навчання;
- складність вибору відповідної швидкості навчання;
- неможливість визначення точок локального і глобального мінімуму, оскільки градієнтний метод їх не розрізняє;
- вплив випадкової ініціалізації вагових коефіцієнтів нейронної мережі на пошук мінімуму функції середньоквадратичної похибки.

Використання нейронних мереж для обробки фінансово-економічних даних підприємств або компаній інколи є недостатньо гнучким. Наприклад, в процесі самонавчання нейромережі не допускається додавання нових нейронів. Складність використання АНМ також обумовлена тим, що розмірність вихідних параметрів має бути визначена до початку навчання. У таких випадках метод АНМ доцільно доповнити генетичними алгоритмами.

Прогнозування з використанням інструментарію генетичних алгоритмів вперше було запропоновано в 70-х роках ХХ століття [62, 80–82]. У другій половині 1980-х рр. до цієї ідеї повернулися у зв'язку з навчанням нейронних мереж. Вони дозволяють вирішувати завдання прогнозування (останнім часом найширше генетичні алгоритми навчання використовуються для банківських прогнозів), класифікації, пошуку оптимальних варіантів, і абсолютно незамінні в тих випадках, коли в звичайних умовах рішення задачі засноване на інтуїції або досвіді, а не на строгому (у математичному сенсі) її описі. Використання механізмів генетичної еволюції для навчання нейронних мереж здається природним, оскільки моделі нейронних мереж розробляються за аналогією з мозком і реалізують деякі його особливості, що з'явилися в результаті біологічної еволюції [83].

Однією з особливостей генетичних алгоритмів є складність

інтерпретації розв'язку та програмна реалізація. Проте, перевагою є ефективність у пошуку глобальних мінімумів багатоекстремальних цільових функцій, оскільки при навчанні нейронних мереж досліджуються великі області допустимих значень вектору параметрів. Генетичні алгоритми дають можливість оперувати дискретними значеннями параметрів нейронних мереж, що може привести до скорочення загального часу навчання [67–69].

Методи індуктивного моделювання у задачах прогнозування.

Особливо відзначимо, що в 70-і роки увагу фахівців в області прогнозування привертала методи індуктивного моделювання. [84–86]. На відміну від дедуктивного методу пізнання, основна ідея індуктивного моделювання полягає у побудові моделей реальних процесів на основі часових рядів, а не на основі теоретичних знань про закономірності функціонування об'єкту.

Ідея методу групового урахування аргументів запропонована академіком О.Г. Івахненком [87], полягає у перекладенні на ЕОМ задачі відшукування моделі, яка найкраще описує процес, що досліджується. Пошук організовано, базуючись на еволюційному принципі, починаючи з найпростіших моделей, які ґрунтуються на опорній функції, що найчастіше має 2 аргументи, наприклад:

$$\begin{aligned}
 y &= a_0 + a_1 x_i x_j, \\
 y &= a_0 + a_1 x_i + a_2 x_j, \\
 y &= a_0 + a_1 x_i + a_2 x_j + a_3 x_i x_j, \\
 y &= a_0 + a_1 x_i + a_2 x_j + a_3 x_i^2 + a_4 x_j^2 + a_5 x_i x_j.
 \end{aligned} \tag{1.14}$$

У загальному випадку, використовується поліном Колмогорова-Габора:

$$y = a_0 + \sum_{i=1}^n a_i x_i + \sum_{i=1}^n a_{ij} x_i x_j + \sum_{i < j < k} a_{ijk} x_i x_j x_k. \tag{1.15}$$

Відомо, що при збільшенні порядку полінома точність наближення ним множини експериментальних даних спочатку зростає, а потім спадає. Тому алгоритм МГУА спрямований на пошук моделі оптимальної складності. Розглянемо більш детально алгоритм МГУА.

Усі доступні дані $\{x_1(t_i)\}, \{x_2(t_i)\} \cdots \{x_m(t_i)\}$ та $\{y(t_i)\}, i = 1, \dots, N$ поділяють на навчальну та контрольну вибірки. Необхідно побудувати модель процесу наступного виду:

$$y(t) = F(x_1(t), x_2(t), \dots, x_m(t)). \quad (1.16)$$

Пошук моделі оптимальної складності виконується у декілька етапів селекції. На кожному етапі перебираються усі пари вхідних параметрів $x_i(t)$ та $x_j(t)$ та будується модель вихідної величини $y(t + \Delta t) = f(x_i(t), x_j(t))$, де Δt – період упередження. У якості функції f вибирається одна з опорних функцій, наприклад, задані рівняннями (1.14) або (1.15). Якщо вибрати внутрішнім критерієм якості моделей критерій мінімальності середньоквадратичного відхилення прогнозів $f(x_i(t_k), x_j(t_k))$ від фактичних майбутніх значень $y(t_{k+1})$, то параметри опорної функції оцінюються методом найменших квадратів, хоча для більш складних опорних функцій можливий вибір іншого критерію та використання нелінійних методів оптимізації для оцінювання параметрів.

Множина отриманих моделей оцінюються згідно вибраного зовнішнього критерію. Таким критерієм може бути мінімальність середньоквадратичного відхилення прогнозів від фактичних майбутніх значень $y(t_{k+1})$ для даних контрольної вибірки. У наступний етап селекції проходять моделі з найкращими показниками зовнішнього критерію. Кількість моделей, які залишаються, задається дослідником і дозволяють обмежити кількість моделей для наступного перебору.

На наступному етапі у якості вхідних величин використовуються відгуки моделей $y_k(x_i, x_j)$, вибраних на попередньому етапі селекції.

Будуються моделі виду $z(t) = f(y_i(t), y_j(t))$. Найкращі за зовнішнім критерієм моделі переходять у наступний етап і т.д. Критерієм завершення алгоритму МГУА є досягнення найкращої точності моделі, після чого точність почне падати.

У монографії [88] запропоновано нечіткий МГУА. У роботах [89, 90] розглянута лінійна інтервальна модель регресії

$$Y = A_0x_0 + A_1z_x + \dots + A_nx_n, \quad (1.17)$$

де $A_i = (a_i, c_i)$ – нечіткі числа трикутного виду, де a_i – центр інтервалу, c_i – його ширина. $c_i \geq 0$, $x_1, x_2, \dots, x_n$ – вхідні змінні ($x_i \in \mathbb{R}$). Тоді відгук Y моделі (1.17) буде теж нечітким числом з центром інтервалу

$$a_y = \sum a_i z_i = a^T \cdot z, \quad (1.18)$$

та шириною інтервалу

$$c_y = \sum c_i |z_i| = c^T \cdot |z|. \quad (1.19)$$

Щоб дана модель коректно описувала навчальну вибірку $\{(y_i, x_1^{(i)}, x_2^{(i)}, \dots, x_n^{(i)})\}$, $i = 1..n$, необхідно знайти такі значення нечітких параметрів $A_i = (a_i, c_i)$ моделі регресії (1.17), щоб межі нечітких відгуків моделі $Y(x_1^{(i)}, x_2^{(i)}, \dots, x_n^{(i)})$ включали реальні значення y_i та б) сумарна ширина нечітких відгуків моделі була б мінімальною. Таким чином, задача нечіткої лінійної регресії зводиться до задачі лінійного програмування:

$$\begin{aligned} \min & (c_0 \cdot n + c_1 \sum_{k=1}^n |x_i^{(k)}| + c_2 \sum_{k=1}^n |x_j^{(k)}| + c_3 \sum_{k=1}^n |x_j^{(k)} \cdot x_j^{(k)}| + \\ & + c_4 \sum_{k=1}^n |x_i^{(k)}|^2 + c_5 \sum_{k=1}^n |x_j^{(k)}|^2), \end{aligned} \quad (1.20)$$

при обмеженнях:

$$\begin{aligned}
 & a_0 + a_1 x_i^{(k)} + a_2 x_j^{(k)} + a_3 x_i^{(k)} \cdot x_j^{(k)} + a_4 \left(x_j^{(k)} \right)^2 + (a_5 x_{kj})^2 - \\
 & - \left(c_0 + c_1 |x_i^{(k)}| + c_2 |x_j^{(k)}| + c_3 |x_i^{(k)} \cdot x_{kj}| + a_4 | \left(x_j^{(k)} \right)^2 | + a_5 | \left(x_j^{(k)} \right)^2 | \right) \leq y_k \\
 & - \left(c_0 + c_1 |x_i^{(k)}| + c_2 |x_j^{(k)}| + c_3 |x_i^{(k)} \cdot x_{kj}| + a_4 | \left(x_j^{(k)} \right)^2 | + a_5 | \left(x_j^{(k)} \right)^2 | \right) \leq y_k \\
 & - \left(c_0 + c_1 |x_i^{(k)}| + c_2 |x_j^{(k)}| + c_3 |x_i^{(k)} \cdot x_{kj}| + a_4 | \left(x_j^{(k)} \right)^2 | + a_5 | \left(x_j^{(k)} \right)^2 | \right) \leq y_k \quad , \\
 & a_0 + a_1 x_i^{(k)} + a_2 x_j^{(k)} + a_3 x_i^{(k)} \cdot x_{kj} + a_4 \left(x_j^{(k)} \right)^2 + a_5 \left(x_j^{(k)} \right)^2 + \\
 & + \left(c_0 + c_1 |x_i^{(k)}| + c_2 |x_j^{(k)}| + c_3 |x_i^{(k)} \cdot x_j^{(k)}| + a_4 | \left(x_j^{(k)} \right)^2 | + a_5 | \left(x_j^{(k)} \right)^2 | \right) \geq y_k
 \end{aligned} \tag{1.21}$$

де k – номер точки вхідних даних $k = 1, 2, \dots, n$. Оскільки для застосування симплекс-методу необхідно мати вимогу невід’ємності шуканих параметрів, то пропонується перейти до двоїстої задачі, розв’язавши двоїсту задачу, відновити значення шуканих параметрів.

У монографії [88] розглянуто методи нечіткої регресії типу (1.17) для нечітких чисел гаусового та дзвоноподібного виду, а також для регресії у вигляді ортогональних поліномів Чебишева та Лагерра та для тригонометричних поліномів. Зведення регресії до нечіткого виду дозволяє подолати відомі проблеми регресійного аналізу, як мультиколінеарність та гетероскедастичність. На увагу заслуговує метод МГУА для нечітких входів та нечіткого виходу, так як дані про ціни сучасних фінансових ринків мають невизначеність і задані інтервальним способом.

1.5. Випадкові процеси та методи побудови марківських моделей при прогнозуванні.

У даному розділі розглядаються моделі для прогнозування фінансових рядів, які базуються на ланцюгах Маркова. На відміну від аналогічних досліджень з фінансового прогнозування [91], не будуються моделі регресійних рівнянь, а пропонується модель машинного навчання (приховані марківські моделі, ПММ) для аналізу фінансових часових

рядів. Також проводиться аналіз попередніх робіт, розкриваються недоліки розроблених раніше моделей та пропонуються шляхи удосконалення, зокрема, використання неперервних сумішей функцій розподілу Гауса та відповідна модифікація класичного алгоритму максимізації правдоподібності. В роботі [92] пропонується подвійно зважений алгоритм максимізації правдоподібності, основна мета якого полягає у подоланні однієї з проблем класичних прихованих марківських моделей – однакової значущості всіх фрагментів даних для прогнозу. Оскільки при прогнозуванні фінансових рядів нові фрагменти змін ціни мають більшу значущість, така модифікація алгоритму навчання моделі дозволяє покращити прогнозні можливості саме для фінансово-економічних часових рядів.

Розглянемо три основні проблеми задач прогнозування часових рядів, які пропонується вирішити за допомогою моделей ПММ. По-перше, характерні властивості відомих фінансових рядів є динамічними, тобто не існує єдиної моделі, яка б працювала весь час. По-друге, важко визначитись між довготривалим трендом та короткочасними флуктуаціями під час торгівлі. Іншими словами, ефективна система повинна уточнювати чутливість з часом. По-третє, важко визначити корисність інформації. Інформація, яка вводить в оману, повинна бути ідентифікована та вилучена.

Робота [92] націлена на подолання вищезазначених проблем за допомогою системи передбачення, яка базується на прихованих марківських моделях (Hidden Markov Model, НММ, далі - ПММ) [8, 93–95]. Модель ПММ включає суміш функцій розподілу для кожного стану, як генератор прогнозу. Суміші функцій Гауса мають бажані властивості, завдяки чому можливо моделювати часові ряди, розподіл яких не відповідає єдиному розподілу Гауса. Потім параметри прихованої марківської моделі доповнюються при кожній ітерації за допомогою

алгоритму максимізації правдоподібності (МП). В кожний дискретний момент часу параметри суміші розподілів, вага кожного компоненту розподілу на виході, та матриці перетворення динамічно оновлюються для забезпечення нестационарності результуючого часового ряду.

Новий експоненційно зважений алгоритм МП дає можливість подолати другу та третю з вищезазначених проблем. Даний алгоритм надає більшого значення новим записам у навчальній вибірці у порівнянні зі старими, налаштовуючи вагові коефіцієнти на зміну волатильності. Баланс між довгими та короткими рядами знайдено в процесі комбінованої торгівлі та передбачення.

Базуючись на експоненціально зваженому алгоритму МП, пропонується подвійно зважений алгоритм МП. Даний алгоритм здатний подолати недолік надчутливості експоненціально зваженого методу максимізації правдоподібності. Показано, що зважений алгоритм МП може бути записаний у формі експоненціальних ковзних середніх від змінних моделі ПММ. Це дозволяє отримати перевагу над існуючими економетричними методами у генеруванні рівнянь переоцінки, основаних на максимізації правдоподібності. Використовуючи технології з теореми про МП, доведено збіжність запропонованих алгоритмів [92].

Експериментальні результати показують, що модель ПММ стабільно показує кращі результати, ніж прибутковість індексу змішаних фондів протягом п'яти 400-денних тестових періодів на проміжку від 1994 до 2002 року. Запропонована модель також стабільно показує кращі результати, ніж відомі алгоритми прогнозування для індексу S&P 500 за показником Шарпа.

Почнемо з визначення властивостей марківського процесу, після чого розглянемо характеристики прихованих марківських моделей. В роботах Рабінера [93] наведено детальний огляд марківських процесів та прихованих марківських моделей та відображені основні проблеми

прихованих марківських моделей. Дані роботи є основоположними при побудові прихованих марківських моделей для прогнозування фінансових часових рядів.

Марківські моделі використовуються для задач аналізу даних, таких як розпізнавання мови, дослідження варіацій температури, біологічних часових рядів та ін. У марківській моделі кожна послідовність спостережень залежить від попередніх елементів у даній послідовності. Розглянемо систему, в якій задано множину станів $S = \{1, 2, \dots, s\}$. У кожний дискретний проміжок часу t , система здійснює перехід з одного стану в інший, згідно із заданими ймовірностями переходів. Позначимо стан системи в момент часу t як s_t .

У багатьох випадках прогноз наступного стану та зв'язана з ним послідовність спостережень залежить тільки від теперішнього стану, це означає те, що ймовірності переходів між станами не залежать від усієї передісторії системи. Такий процес називається марківським процесом першого порядку.

Властивості таких процесів можна описати наступними рівняннями: 1) обмежений горизонт

$$P(X_{t+1} = s_k | X_1, \dots, X_t) = P(X_{t+1} = s_k | X_t) \quad (1.22)$$

2) часова інваріантність

$$= P(X_2 = s_k | X_1) \quad (1.23)$$

Оскільки переходи між станами інваріантні у часі, можна записати матрицю ймовірностей переходів A з елементами:

$$a_{ij} = P(X_{t+1} = s_j | X_t = s_i) \quad (1.24)$$

a_{ij} є ймовірностями, отже: $a_{ij} \geq 0$;

$$\forall i, j \sum_{j=1}^N a_{ij} = 1 \quad (1.25)$$

Також необхідно зафіксувати ймовірність кожного стану на початку певного часового інтервалу та початковий розподіл ймовірностей:

$$\pi_i = P(X_1 = s_i) \quad (1.26)$$

при

$$\sum_{i=1}^N \pi_i = 1. \quad (1.27)$$

У явній марківській моделі стани, між якими здійснюються переходи та функції ймовірностей є відомими, тому можна розглядати послідовність станів як вихідні дані.

Ланцюги Маркова можна застосувати для моделювання рухів ціни акцій, ввівши множину дискретних станів системи для моделювання зміни ціни певного виду акцій. Ланцюг Маркова можна зобразити у вигляді графу переходів між станами. Стани в ланцюгу Маркова можна зв'язати зі зміною ціни: “значне збільшення”, “невелике збільшення”, “незмінність”, “невеликий спад”, “значний спад”.

Марківська модель, яка описана вище, є обмеженою, тому ми розширюємо її для покращення можливостей в приховану марківську модель (далі ПММ). В ПММ процес, який генерує послідовність, що розглядається, є невідомим. Кількість станів, ймовірності переходів та початковий стан також невідомі.

Замість квантування кожної зміни детермінованого виходу набором дискретних станів (таких як велике зростання, незначний спад тощо), кожний стан прихованої марківської моделі зв'язаний з ймовірнісною функцією. В момент часу t , послідовність станів o_t генерується ймовірнісною функцією $b_j(o_t)$, яка зв'язана зі станом j , та має такий вигляд:

$$b_j(o_t) = P(o_t | X_t = j). \quad (1.28)$$

Розглянемо приклад прихованої марківської моделі, яка описує діяльність інвестора з передбачення змін ринкової ціни, оснований на множині стратегій інвестування. Припустимо, що існує 5 стратегій інвестування, кожна з яких дає відмінний прогноз ринкової ціни на наступний день (великий спад, невеликий спад, без змін, невелике збільшення, значне збільшення). Відомо, що інвестор здійснює прогноз кожен день лише згідно однієї стратегії.

Яким чином професіонал здійснює прогноз – невідомо. Також невідомо, якою стратегією він буде користуватись кожного наступного дня.

Для відповіді на перше запитання створимо приховану марківську модель як модель фінансового експерта. Кожна з п'яти стратегій відповідає прихованому стану в ПММ. Множина змін ціни (явних станів) містить усі символи (коди) змін ціни, та кожна зміна зв'язана з певним станом. Зі стратегіями зв'язані різні розподіли ймовірностей символів змін ціни. Кожний стан зв'язаний з іншими станами за допомогою функції ймовірностей переходів.

На відміну від марківських ланцюгів, прихована марківська модель передбачає вибір трейдером окремої стратегії після використання в минулий раз однієї з інших, тільки залежно від ймовірності послідовності спостережень. Шлях до удосконалення ефективності ПММ полягає у тому, що вибирається найкраща послідовність стратегій на основі наявної послідовності станів. Вводячи функції розподілу ймовірностей для станів ПММ, збільшуємо можливості моделі у презентабельності, порівняно з явною моделлю, при якій зв'язки “стратегія-стан” є фіксованими.

Алгоритм генерації послідовності прогнозних станів включає наступні кроки:

1. Вибрати початковий стан (початкову стратегію) $X_1 = i$, згідно з початковим розподілом станів Π .
2. Задати номер проміжку часу $t = 1$, з якого відраховується послідовність станів.
3. Отримати прогноз, згідно зі стратегією даного стану i , тобто $b_{ij}(o_t)$. Таким чином, ми маємо прогноз для кожного наступного дня.
4. Здійснити перехід до нового стану X_{t+1} , згідно розподілу ймовірностей переходів між станами.
5. Задати час $t = t + 1$ та перейти до кроку 3, якщо $t \leq T$, інакше завершити роботу.

Очевидно, що існує реальна відповідність між аналізом послідовностей даних та ПММ. Особливо у випадку аналізу фінансових часових рядів та прогнозування, коли дані можуть бути далекі від генерованих певним стохастичним процесом, який невідомий публіці. Аналогічно вищезазначеному прикладу, якщо кількість стратегій професіонала та їх суть невідомі, можна взяти уявного професіонального інвестора і вважати його зовнішньою силою, яка рухає ціну вгору чи вниз. ПММ має більш широкі можливості, ніж звичайні ланцюги Маркова. Ймовірності переходів, як і функції розподілу ймовірностей появи послідовності станів, уточнюються в процесі “навчання” алгоритму.

Крім схожості властивостей ПММ та даних часового ряду, алгоритм максимізації правдоподібності (МП) представляє собою клас ефективних методів “навчання” моделі. Маючи значні об’єми даних, які є результатом деякого прихованого процесу, можна створити ПММ, а алгоритм МП дає можливість оцінити параметри моделі, що найкраще відповідають відомим даним торгівлі.

Загальний підхід ПММ включає технологію некерованого навчання, яка дозволяє знайти нові характерні риси досліджуваного процесу. Алгоритм навчання здатний приймати вхідні послідовності різної

довжини, без необхідності зводити навчання до “шаблону”. Удосконалений авторами роботи [92] алгоритм МП дає можливість збільшити “вагу” більш свіжих даних при навчанні, а також відкривається можливість налаштовувати баланс між вимогами до чутливості та точності.

ПММ застосовуються під час розпізнання мови, у біоінформатиці та аналізі різних часових рядів, таких як дані погоди, дефекти напівпровідників та ін. Вайгенд та Ши вперше застосували ПММ в аналізі фінансових часових рядів.

Можна виділити три фундаментальних задачі ПММ, які є важливими для прихованих марківських моделей:

1. Для відомої моделі $\mu = (A, B, \Pi)$, та послідовності явних станів $O = (o_1, \dots, o_T)$, обчислити ймовірність певної послідовності явних станів моделі. Тобто обчислити ймовірність $P(o_j)$.

2. Для відомої моделі та послідовності явних станів O , визначення прихованої послідовності станів $(1, \dots, N)$, що найкраще “пояснює” дану послідовність явних станів.

3. Для відомої послідовності O та простору можливих моделей визначити параметри та знайти модель, яка максимізує ймовірність $P(O_j)$.

Перша задача може використовуватись для визначення, які з “навчених” моделей найбільш адекватно описують послідовність явних станів, яка спостерігається. У випадку фінансових часових рядів для відомої історії змін ціни або індексу, яка модель найкраще пояснює дану історію. Друга задача полягає в знаходженні прихованого шляху. Звернувшись до вищенаведеного прикладу зі стратегіями, задача полягає у визначенні послідовності стратегій, які щоденно використовує інвестор-професіонал. Третя задача є найбільш цікавою, так як стосується “навчання” моделі. В аналізі реальних фінансових часових рядів, розглядається навчання моделі на основі історії минулих рухів ціни. Тільки

після “навчання” ми можемо використовувати модель для подальшого прогнозування.

В загальному вигляді ПММ є кортеж (S, K, Π, A, B) , де:

1. $S = \{1, \dots, N\}$ – множина станів. Стан системи в момент часу t позначимо s_t .

2. $K = \{k_1, \dots, k_M\}$ – вихідний алфавіт. У випадку дискретних станів, M є кількістю дискретних станів. У вищезгаданому прикладі M дорівнює кількості можливих змін ціни.

3. Початковий розподіл станів $= \{\pi_i\}$, де $i \in S$, π_i визначається як

$$\pi_i = P(s_1 = i). \quad (1.29)$$

4. Розподіл ймовірностей переходів між станами $A = \{a_{ij}\}; i, j \in S$.

$$a_{ij} = P(s_{t+1} | s_t); 1 \leq i, j \leq N \quad (1.30)$$

5. Розподіл ймовірностей явних станів $B = b_j(o_t)$. Функція ймовірності для кожного стану задається наступним чином:

$$b_j(o_t) = P(o_t | s_t = j) \quad (1.31)$$

Після формалізації задачі у вигляді ПММ та припущення, що модельовані дані, згенеровані за допомогою ПММ, необхідно обчислити ймовірності послідовностей станів системи та можливу послідовність прихованих станів. Також проводиться “навчання” параметрів моделі, яке базується на відомих даних та отримується більш точна модель. Далі “навчену” модель можна використовувати для передбачення нових даних. Алгоритм симуляції прихованої марківської моделі включає наступні кроки:

Виберемо початковий стан x_1 по розподілу π .

Для усіх t від 1 до T :

а) Виберемо спостережувану d_t за розподілом

$$b_j(k) = p(v_k | x_j).$$

б) Виберемо наступний стан за розподілом

$$a_{ij} = p(q_{t+1} = x_j | q_t = x_i).$$

Алгоритм знаходження ймовірності послідовності спостережень у даній моделі:

1. Ініціалізація: $\alpha_1(i) = \pi_i b_i(d_1)$.
2. Крок індукції: $\alpha_{t+1}(j) = \left[\sum_{i=1}^n \alpha_t(i) a_{ij} \right] b_j(d_{t+1})$.
3. Після T ітерацій підрахуємо приховану ймовірність:

$$p(D | \lambda) = \sum_{i=1}^n \alpha_T(i).$$

Алгоритм знаходження найбільш ймовірних послідовностей станів:

1. Ініціалізація: $\delta_1(i) = \pi_i b_i(d_1)$, $\psi_1(i) = \{ \}$.
2. Крок індукції:

$$\delta_t(j) = \max_{1 \leq i \leq n} [\delta_{t-1}(i) a_{ij}] b_j(d_t),$$

$$\psi_t(j) = \arg \max_{1 \leq i \leq n} [\delta_{t-1}(i) a_{ij}].$$

3. Після T ітерації визначити:

$$p^* = \max_{1 \leq i \leq n} \delta_T(i), \quad q_T^* = \arg \max_{1 \leq i \leq n} \delta_T(i).$$

4. Визначити послідовність: $q_t^* = \psi_{t+1}(q_{t+1}^*)$.

Розглянутий ітераційний алгоритм дозволяє уточнити параметри моделі та послідовність прихованих станів для даного набору спостережень.

У даному підрозділі теоретично описується загальна форма ПММ, та матеріал ілюструється за допомогою прикладів. Формалізується марківська властивість, яка є основою даних моделей. Далі розглядаються три основні задачі, зв'язані з ПММ. Розглянута загальна структура моделі та введені розширення даної базової моделі, описані технології розв'язання поставлених задач. Для організації “навчання” системи передбачення використовуються алгоритми Вітербі та Баум-Велча.

Проаналізуємо попередній досвід застосування прихованих марківських моделей для прогнозування часових рядів. Першою спробою застосувати ПММ для прогнозування в фінансах здійснили вчені А. Вайгенд та Ш. Ши [95]. Ними була запропонована модель експертів для прогнозування динаміки індексу S&P 500. Під експертами, які були закодовані дискретними станами ПММ мались на увазі три нейронних мережі з різними налаштуваннями. Перша нейромережа була налаштована на прогнозування даних з невеликою волатильністю, друга була навчена на високоволатильних даних, а третя відповідала за “викиди”. Під “викидами” мається на увазі частина даних, які не увійшли до двох вищезазначених груп.

Під час прогнозування створювався торговий пул, який являв собою середовище з обмеженими у часі даними про динаміку ціни. Даний пул оновлювався за допомогою додавання свіжих даних про ціну та вилучення з пулу старих. Моделі в даному пулі проходили “навчання” з ціллю генерації прогнозу на наступний торговий день. Навчання експерта, оснований на ПММ являло собою коригування значень матриці ймовірностей переходів A .

Модель Вайгенда та Ши має наступні недоліки. По-перше, в торговому пулі значення кожного фрагменту були однаковими, що при великій кількості даних призводить до проблеми перенавчання. Для вирішення цієї проблеми запропоновано експоненціально зважений алгоритм МП, у якому

для свіжих даних надається більша значущість. Наступною проблемою є те, що при “навчанні” для трьох моделей використовується однаковий алгоритм оцінювання коефіцієнтів, не враховуючи специфічні можливості кожного виду моделі. Навіть якщо класифікувати фрагменти даних за показником волатильності, то такий поділ буде суб’єктивним, а ряди з проміжними значеннями волатильності не обов’язково демонструватимуть однакові властивості. В роботі [92] наводяться моделі з сумішами розподілів, в яких налаштовуються окремі параметри кожного компоненту суміші.

У попередніх підрозділах формули для адаптації параметрів моделі були викладені для випадку дискретної послідовності спостережень. Але зміни ціни не є дискретними значеннями, процедура квантифікації неперервних значень дискретизованими станами призводить до втрати інформації. Тому пропонується використовувати суміші функцій розподілу Гауса як модель неперервних спостережень.

У випадку неперервної функції розподілу ПММ, функція $b_j(o_t)$ є неперервною функцією розподілу або суміш неперервних функцій розподілу:

$$b_j(o_t) = \sum_{k=1}^M w_{jk} b_{jk}(o_t), j = 1, \dots, N, \quad (1.32)$$

де M – кількість сумішей та w є вага кожної суміші. Ваги сумішей мають наступні обмеження:

$$\sum_{k=1}^M w_{jk} = 1, j = 1, \dots, N; w_{jk} \geq 0; j = 1, \dots, N; k = 1, \dots, M. \quad (1.33)$$

Кожна функція $b_{jk}(o_t)$ є D -вимірною логарифмічно ввігнутою функцією розподілу з середнім вектором μ_{jk} та матрицею коваріацій Σ_{jk} :

$$b_{jk}(o_t) = N(o_t, \mu_{jk}, \Sigma_{jk}). \quad (1.34)$$

Основна задача “навчання” – визначення параметрів складових функцій розподілу ймовірності. Таке D-вимірне навчання вимагає вибору найбільш корельованих часових рядів у якості вхідних даних. Наприклад, для прогнозування ціни на акції IBM, є сенс включити у вхідні дані ціни акцій Microsoft та інших високотехнологічних компаній.

Звичайна функція розподілу Гауса має наступний вигляд:

$$b_{jk}(o_t) = N(o_t, \mu_{jk}, \Sigma_{jk}) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(x - \mu_{jk})^2}{2\Sigma_{jk}^2}\right), \quad (1.35)$$

де o_t – дійсне число. Багатовимірна функція розподілу Гауса запишеться так:

$$\begin{aligned} b_{jk}(o_t) &= N(o_t, \mu_{jk}, \Sigma_{jk}) = \\ &= \frac{1}{(2\pi)^{\frac{D}{2}} |\Sigma_{jk}|^{\frac{1}{2}}} \exp\left(-\frac{1}{2} (o_t - \mu_{jk})^T \Sigma_{jk}^{-1} (o_t - \mu_{jk})\right), \end{aligned} \quad (1.36)$$

де o_t – вектор явних значень прогнозованого ряду. Пропонується модифікація алгоритму Баума-Велча [92] для випадку багатовимірних часових рядів з дійсними значеннями. Вводяться наступні допоміжні змінні:

$$\delta_i(t) = \sum_{m=1}^M \delta_{im}(t) = \sum_{m=1}^M \frac{1}{P} a_i(t-1) a_{ji} w_{im} b_{im}(o_i) \beta_i(t). \quad (1.37)$$

Далі необхідна змінна для позначення ймовірності знаходження у стані j в момент часу t з k -м компонентом суміші, яка пояснює даний набір спостережень o_t :

$$\nu_{j,k}(t) = \frac{\alpha_j(t) \beta_j(t)}{\sum_{j=1}^N \alpha_j(t) \beta_j(t)} \cdot \frac{w_{jk} N(o_i, \mu_{jk}, \sigma_{jk})}{\sum_{k=1}^M w_{jk} N(o_i, \mu_{jk}, \sigma_{jk})} \quad (1.38)$$

З новими допоміжними змінними запишемо модифіковані формули уточнення параметрів:

$$\mu'_{im} = \frac{\sum_{t=1}^T \delta_{im}(t) o_t}{\sum_{t=1}^T \delta_{im}(t)} \quad (1.39)$$

$$\sigma'_{im} = \frac{\sum_{t=1}^T \delta_{im}(t) (o_t - \mu'_{im}) (o_t - \mu'_{im})'}{\sum_{t=1}^T \delta_{im}(t)}, \quad (1.40)$$

$$w'_{im} = \frac{\sum_{t=1}^T \delta_{im}(t) o_t}{\sum_{t=1}^T \delta_i(t)}, \quad (1.41)$$

$$a'_{ij} = \frac{\sum_{t=1}^T a_i(t) \alpha_i(t) a_{ij} b_j(o_{t+1}) \beta_i(t+1)}{\sum_{t=1}^T \alpha_i(t) \beta_i(t)}. \quad (1.42)$$

В даному розділі розглянуто ПММ, адаптовану для прогнозування фінансових часових рядів, зокрема для індексу S&P 500. Розглянуті питання про вибір функції щільності розподілу приростів досліджуваної величини, правило коригування параметрів моделі, яке базується на максимізації правдоподібності. Кожен стан ПММ зв'язаний з сумішню функцій розподілу Гауса з різними параметрами, які дають кожному стану різні прогнозні переваги. Деякі стани більш чутливі до низьковолатильних даних, інші більш точні на високоволатильних проміжках часу. Модернізований алгоритм МП приділяє більше уваги свіжим даним. Індекс Dow Jones включений в експерименти як допоміжний індикатор для визначення волатильності ряду S&P 500.

Зважений алгоритм максимізації правдоподібності для ПММ

Розглянемо більш детально традиційний алгоритм МП з точки зору конструювання функції ймовірності та Q -функції. У попередніх роботах не приділяється достатньої уваги випадку затухання значення фрагменту даних з часом. Прийнято, що кожен фрагмент даних у послідовності навчання має однакову важливість та остаточне оцінювання базується на

припущенні цієї рівної важливості, не звертаючи увагу, наскільки велика вибірка даних для навчання та проміжок часу, який описують дані. Такий підхід є ефективним при невеликих фрагментах навчання та нечутливості, або розмірність часу не є важливою. Традиційні моделі ПММ, основані на максимізації правдоподібності, показують високі результати в обробці та розпізнаванні мови та відео, де послідовність фрагментів даних не має великого значення.

Але фінансові часові ряди відрізняються від вищезгаданих сигналів своїми унікальними характеристиками. Необхідно більше звернути увагу на нові дані, залишаючи у вибірці навчання і старі, віддалені у часі фрагменти, але використовуючи їх зі зменшеною довірою. Усвідомлюючи, що зміна важливості даних з часом та мотивована цим концепція експоненційного ковзного середнього (далі ЕКС) фінансових часових рядів, в роботі [92] пропонується модернізація експоненційно зваженого алгоритму МП. Показано, що остаточне перевизначення параметрів ПММ та функцій розподілу може бути виражене у формі комбінації експоненційних ковзних середніх (ЕКС) для змінних, що містять проміжні стани. Для знаходження кращого балансу між короткочасними та довготривалими трендами, авторами [92] пропонується подвійно зважений алгоритм МП, який має властивості налаштовувати чутливість за допомогою включення інформації про волатильність.

Доведено збіжність експоненційно зваженого алгоритму МП за допомогою зваженої версії теореми про максимізацію правдоподібності.

Нехай функція розподілу ймовірності $p(X|\Theta)$, яка показує ймовірність послідовності спостережень $X = \{x_1, \dots, x_N\}$ для даного параметру Θ . В моделі прихованих марківських гаусових сумішей (далі ПМГС) Θ включає множину середніх та коваріацій гаусових сумішей, зв'язаних з кожним станом, разом з матрицею ймовірностей переходу та ваговими коефіцієнтами кожного компоненту суміші.

Вважаючи, що дані є незалежними та рівномірно розподіленими за законом p , ймовірність спостерігати всю послідовність даних обчислюється наступним чином:

$$P(X|\Omega) = \prod_{i=1}^N p(x_i|\Omega) = L(\Omega|X) \quad (1.43)$$

Функція L називається функцією правдоподібності для параметрів при заданому наборі даних. Ціль алгоритму МП полягає в знаходженні вектору параметрів Θ^* , який відповідає максимальному значенню L :

$$\Omega^* = \arg \max_{\Omega} L(\Omega|X). \quad (1.44)$$

Далі алгоритм МП знаходить математичне сподівання для логарифмічної функції правдоподібності $\log p(X; Y|\Theta)$ відносно невідомих значень Y для даної явної частини X та параметрів останньої оцінки. Виберемо Q -функцію в якості позначення для цього очікування:

$$\begin{aligned} Q(\Theta, \Theta^{(i-1)}) &= E [\log p(X, Y|\Theta) | X, \Theta^{(i-1)}] = \\ &= \int_{y \in \gamma} \log p(X, y|\Theta) f(y|X, \Theta^{(i-1)}) dy. \end{aligned} \quad (1.45)$$

де $\Theta^{(i-1)}$ є оцінки з минулої ітерації алгоритму МП. Базуючись на $\Theta^{(i-1)}$ намагаємось знайти нові оцінки Θ , які оптимізують функцію Q . γ є простір можливих значень для змінної y .

При обчисленні функції Q вважається, що кожен фрагмент даних має рівну важливість при оцінці нових параметрів. Цей момент є важливим у задачах розпізнавання мови та візуальних сигналів, але фінансові часові ряди мають характеристики, для яких параметр часу стає дуже важливим. З плином часу, вираженість старих фрагментів затухає, з'являються нові риси часового ряду. Якщо локальна послідовність даних, яка використовується у “навчанні” є достатньо великою, нові риси або тренди, що виникають, неминуче втратять свою чіткість.

Тому необхідно створити новий алгоритм МП, який надає більше важливості для більш свіжих даних у процедурі “навчання”. Для розв’язання цієї задачі, необхідно по-перше перетворити Q -функцію. Нагадуємо, що Q -функція є мірою “правдоподібності” повного набору даних відносно прихованих даних Y для даної послідовності явних станів X та оцінки останніх параметрів Θ . Пропонується включити часово-залежні вагові коефіцієнти η , очікуючи відобразити інформацію про часовий параметр.

Нова форма функції правдоподібності Q^* матиме наступний вигляд:

$$\begin{aligned} Q^*(\Theta, \Theta^{(i-1)}) &= E [\log \eta p(X, Y|\Theta) | X, \Theta^{(i-1)}] = \\ &= \int_{y \in \gamma} \log \eta p(X, y|\Theta) f(y|X, \Theta^{(i-1)}) dy. \end{aligned} \quad (1.46)$$

Аналогічно формулам (1.39)-(1.42) виводяться ітераційні формули уточнення параметрів моделі при зваженому врахуванні фрагментів даних.

Висновки до розділу 1

Короткий огляд підходів і економіко-математичних методів прогнозування часових рядів дозволяє зробити висновки, що не існує одного універсального, задовольняючого всім вимогам, методу прогнозування, що не володіє недоліками. Кожен підхід і кожен метод мають свої переваги, недоліки, межі застосування. У світовій економічній літературі кількість методів прогнозування обчислюється багатьма десятками. Варто відзначити, що ці методи базуються або на кореляційно-регресійних моделях, або на трендах, для представлення яких вибирається найбільш відповідні екстраполяційні залежності.

Величезний досвід математичного моделювання динамічних (еволюційних) процесів, накопичений в світі за останні десятиліття, невимірно розширив і багато в чому змінив сталі уявлення про

адекватність існуючих математичних моделей суті цих процесів. Стало очевидно, що класичного арсеналу математичного моделювання, що базується на лінійній парадигмі (при якій малі збурення вхідних даних системи у незначній мірі змінюють її траєкторію), у багатьох випадках явно недостатньо для побудови адекватних математичних моделей. Ця обставина зумовила фундаментальний перегляд лінійної концепції і перехід на так звану нелінійну парадигму (nonlinear science) в математичному моделюванні, при якій малі збурення вхідних даних або значень змінних динамічної системи можуть в катастрофічно великій мірі змінити її траєкторію через складність самої системи і хаотичність її поведінки). Практична цінність указаної парадигми обумовлена тим, що за допомогою її основних положень стає можливим біоляш адекватно відображати специфічні характеристики ієрархічності, конкретної динаміки і високий рівень невизначеності, властиві реальним соціальним, економічним, фінансовим, фізичним та іншим складним процесам і системам. Перехід на нову концепцію викликав необхідність створення принципово нових інструментальних засобів математичного моделювання, зокрема таких, як фрактальна геометрія, фрактальний аналіз, методи детермінованого хаосу і ін. Через цю обставину для побудови прогнозної моделі запропоновано новий підхід, який базується на використанні інструментарію складних ланцюгів Маркова з використанням ієрархії приростів часу та дискретного Фур'є-продовження для прогнозування низькочастотної складової часового ряду.

РОЗДІЛ 2

МЕТОДИ ПРОГНОЗУВАННЯ НА ОСНОВІ СКЛАДНИХ ЛАНЦЮГІВ МАРКОВА ТА ДИСКРЕТНОГО ФУР'Є-ПРОДОВЖЕННЯ

2.1. Метод прогнозування на основі складних ланцюгів Маркова

Розглянемо поняття ланцюгів Маркова. Припустимо існує послідовність дискретних станів певної системи. З цієї послідовності можна визначити ймовірності переходу з одного стану в інший. Простим ланцюгом Маркова є випадковий процес, в якому ймовірності наступного стану залежать тільки від попереднього стану та не залежать від усіх інших станів. На відміну від простого, складним ланцюгом Маркова називають випадковий процес, в якому ймовірності наступного стану залежать не тільки від наявного, а від послідовності декількох попередніх станів (передісторії) [96, 97]. Основна властивість ланцюгів Маркова задана рівнянням (1.22). Кількість станів в передісторії називається порядком ланцюга Маркова.

Теорія простих ланцюгів Маркова широко викладена в літературі, наприклад в [96, 97], що стосується ланцюгів Маркова вищих порядків, то в літературі даний термін використовується рідко. Наприклад, в статті [98] автор використовує термін “ланцюги Маркова” для процесів з пам'яттю об'ємом декілька попередніх значень. У роботах [99, 100] пропонується використання класу періодичних ланців Маркова для побудови прогнозів часових рядів технічних систем.

Ланцюг Маркова порядку вище першого можна звести до простого

ланцюга Маркова за допомогою узагальнення поняття стану: “наявний стан”, як ряд послідовних станів [49]. В цьому випадку апарат простих ланцюгів Маркова може бути застосований до складних.

Припустимо, що процес, динамічний ряд якого досліджується, породжений певною моделлю детермінованого хаосу, що означає існування причинно-наслідкової залежності наступних станів від передісторії. Для точного оцінювання ймовірностей наступного стану, необхідно врахувати кожен з попередніх, оцінивши ймовірності переходів при даній передісторії. Оскільки знання про причинно-наслідкову залежність обмежена даною реалізацією ряду, то ймовірності можуть бути оцінені тільки наближено. Основна задача сучасних методів прогнозування максимально використати інформацію з відомого відрізка ряду для побудови моделі ряду та використати побудовану модель, виокремивши з можливих сценаріїв продовження ряду найбільш ймовірні.

Ряд вихідних значень, призначений для прогнозування, необхідно перевести в ряд дискретних номерів станів. Необхідно зафіксувати кількість станів s дискретного ланцюга Маркова, кожен з яких зв’язуємо з абсолютною або відносною зміною величини вихідного сигналу за певний проміжок часу Δt . Найпростішим прикладом може бути 2 стани: зростання (стан 1) та спадання (стан 2). В загальному випадку всі можливі прирости вихідного ряду необхідно класифікувати на s класів [4]. Способи розбиття будуть розглянуті у наступних підрозділах даної роботи.

Далі здійснюється прогнозування ряду дискретизованих станів. При заданому порядку ланцюга Маркова та останньому узагальненому стані вибирається найбільш ймовірний стан в якості наступного. У випадках неоднозначності при визначенні найбільш ймовірного стану застосовується алгоритм, який дозволяє зменшити кількість можливих сценаріїв прогнозу. Таким чином, маємо ряд прогнозованих станів, які для відомого

останнього значення ряду можуть бути перетворені на дискретизований ряд прогнозних значень.

Обчислення приростів, прогнозування та послідуєче відновлення можна здійснити для різних приростів часу Δt . Для ефективного використання інформації, представленої в наявному часовому ряді, прогнозування здійснюється для різних приростів часу $\Delta t = 1, 2, 4, 8, \dots$, або більш складної ієрархії приростів, яка розглянута в даній роботі в підрозділі 2.3, та послідовного “склеювання” результатів прогнозування, отриманих при різних інтервалах дискретизації Δt .

Процедура прогнозування та склеювання є ітераційною та проводиться, починаючи з менших приростів, додаючи на кожному кроці прогноз з більшим приростом часу.

Оскільки при збільшенні кроку дискретизації часу Δt зменшується статистика для визначення ланцюгів Маркова, найбільший крок дискретизації, який приймає участь у прогнозуванні, обмежується. Для доповнення прогнозу низькочастотною складовою використовується наближення нульового порядку у вигляді лінійного тренду [49], або комбінації лінійного тренду та гармонійних коливань [14]. Дана процедура розглядається в розділі 2.5 даної роботи.

Стани в складних ланцюгах Маркова та підходи до їх визначення.

Як зазначено вище, стани дискретного ланцюга Маркова визначаються величинами абсолютних або відносних приростів величини, що прогнозується. Досліджуваний процес описується у вигляді часового ряду ціни $p(t)$ із заданим проміжком дискретизації Δt

$$p_{ti} = p(t_0 + i\Delta t). \quad (2.1)$$

Основою для класифікації станів є приріст, або прибутковість ряду.

Розглядається абсолютні (2.2) та відносні (2.3) прирости ряду.

$$r_{abs}(t) = p_t - p_{t-\Delta t} \quad (2.2)$$

$$r_{rel}(t) = \frac{p_t - p_{t-\Delta t}}{0.5 \cdot (p_t + p_{t-\Delta t})}; \quad (2.3)$$

де p_t – вхідний ряд динаміки ціни, Δt – проміжок дискретизації, який вибрано для аналізу. Математичне сподівання ряду прибутковостей дорівнює нулю, а дисперсія є фактично мірою волатильності ряду. Ряд вихідних значень необхідно перетворити у ряд дискретних станів. Позначимо кількість вибраних станів s , кожен з яких пов'язаний із абсолютною ($r_{abs}(t)$) або відносною ($r_{rel}(t)$) прибутковістю сигналу, обчислену з кроком Δt . Алгоритм розбиття на стани відіграє важливу роль для якості прогнозу, оскільки необхідно щоб система станів інформативно описувала вхідний часовий ряд дійсних величин, дозволила виявити закономірності у послідовностях станів, та з іншого боку було використано мінімальна необхідна кількість дискретних станів, при якому кодування не погіршувало б точність представлення часового ряду. У дисертації запропоновано ефективні алгоритми дискретизації ряду прибутковостей у вигляді послідовності дискретних станів, які враховують особливості негаусового розподілу прибутковостей фінансово-економічних часових рядів. Пропонуються наступні способи класифікації прибутковостей у стани: класифікація на основі принципу рівномірності за кількістю представників в класах та на основі принципу рівномірності за відхиленням, а також їх комбінації для різних модулів відхилень. На основі значень прибутковостей r_t здійснюється класифікація та перетворення значень ряду в ряд дискретних станів. Один з принципів проведення класифікації рівномірність за кількістю представників класів. Дана класифікація поділяє множину всіх приростів r_t на s рівних по кількості груп. Обчислені прирости ряду з даною дискретизацією упорядковуємо за зростанням та здійснюємо поділ відсортованого масиву на рівні частини.

Таким чином визначаються граничні значення $r_{lim,i}$, які використовуються при перетворенні прибутковостей на номери класів. Проблему може створити велика кількість однакових станів, що спричинює однакові межі декількох сусідніх станів. Це створює номери станів без жодного представника, що робить необхідним корекцію розбиття з ціллю досягнути найбільш можливої рівномірності в розподілу на стани. Класифікація відбувається за наступним алгоритмом:

$$s_t = (i | r_{lim,i-1} > r_t \geq r_{lim,i}) \quad (2.4)$$

де s_t – номер стану, який відповідає моменту часу t , для якого обчислена прибутковість r_t ; i – номер стану ($1 \dots s$), інтервал $(r_{lim,i-1}, r_{lim,i}]$ якого відповідає обчисленій прибутковості r_t . Крім інтервалу прибутковостей, заданого зазначеними вище граничними значеннями $(r_{lim,i-1}, r_{lim,i}]$, для кожного стану вибирається середнє значення прибутковості $r_{avg,i}$, яке буде використовуватись при відновленні значень ряду за прогнозованими дискретними станами.

Пропонується декілька алгоритмів класифікації відхилень на стани. Алгоритми класифікації виділяють межеві прирости для s класів, чим визначають, до якого класу відноситься той чи інший приріст у вибірці навчання. Розглянемо алгоритми класифікації:

1. Розподіл, рівномірний за кількістю представників кожного класу;
2. Розподіл, рівномірний за величиною відхилень;
3. Комбінований розподіл (межі крайніх лівого та правого стану визначається рівномірними за кількістю, усі інші - рівномірні за відхиленням);
4. Комбінований розподіл при можливості регулювати кількість представників крайніх станів.

Розподіл, рівномірний за кількістю представників, проводиться за наступною схемою. Обчислюємо прирости ряду з даною дискретизацією,

упорядковуємо їх та ділимо відсортований масив на рівні частини. Таким чином, ми досягаємо рівномірності по кількості представників в станах. Проблему може створити велика кількість однакових станів, що спричинює однакові межі декількох сусідніх станів. Це створює номери станів без жодного представника, що робить необхідним корекцію розбиття з ціллю досягнути найбільш можливої рівномірності розподілу частот представників станів.

Другий спосіб розбиття на стани полягає в тому, щоб розбити інтервал від мінімального (від'ємного) до максимального відхилення на рівні по відхиленню частини. В цьому випадку рівномірність по кількості представників в станах не зберігається. Це далеко не всі можливі способи розбиття.

Для кожного стану вибирається середнє значення, яке буде використовуватись при відновленні ряду по прогнозованим станам.

Оскільки реальна причинно-наслідкова залежність в процесі, який представлений динамічним рядом, невідома, то пропонується декілька способів розбиття на стани, при яких, по-перше, переходи між станами були представлені достатньою кількістю прецедентів, а по-друге, усереднення відхилень всередині станів не вплинуло б суттєво на точність одержаного прогнозу.

Для перевірки ефективності розбиття пропонується провести процедуру дискретизації та класифікації приростів на стани на кожному рівні Δt ієрархії приростів, а потім по відомим станам для кожного Δt відновити значення ряду та провести процедуру склеювання. Оскільки ряди станів повністю відповідають вихідному ряду, то отримується крива, відхилення якої від вихідної спричинені лише похибкою усереднення всередині станів (похибка квантування). Таким чином, прийнявши певне значення кількості станів s , та провівши процедуру дискретизації,

відновлення та склеювання (без прогнозування), одержуємо абсолютну похибку дискретизації (квантування).

Розглянемо процес дискретизації та відновлення дискретизованих значень функції $y = \sin(x)$ для $s = 3, 6, 9, 12$ у випадку простої та складної ієрархії.

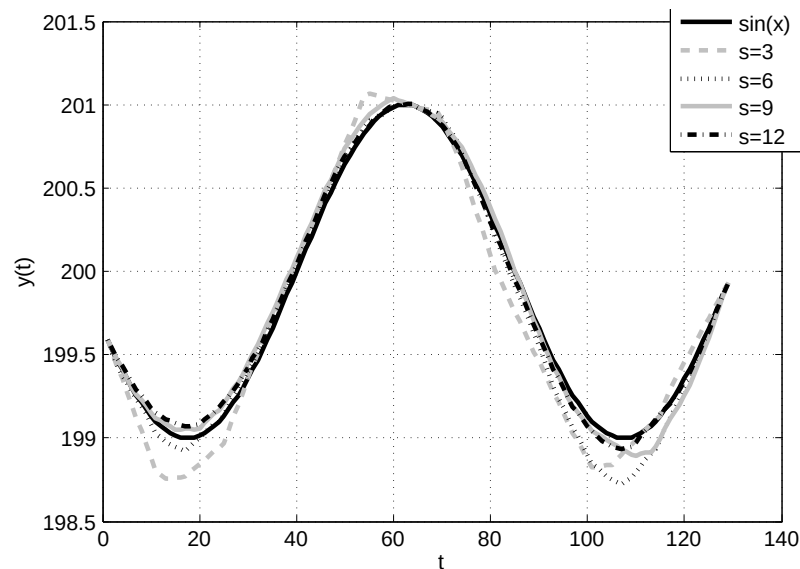


Рис. 2.1. Дискретизація та відновлення $y = \sin(x)$. Проста ієрархія

Із рис. 2.1-2.4 видно, що складна ієрархія дає більш точне відновлення, так як більш густа сітка дискретизації часу виправляє похибки квантування, які виникають на інших рівнях Δt ієрархії приростів часу. Також зі збільшенням кількості станів точність відновлення збільшується, але необхідно мати на увазі, що вибір кількості рівнів квантування обмежений тим, що для прогнозування необхідне визначення ймовірностей переходів з достатньою точністю, для чого необхідна достатня кількість переходів між виокремленими станами.

Алгоритм побудови багатокрокового прогнозу

Розглянемо послідовність операцій, необхідних для побудови прогнозу вихідного ряду. Для цього необхідно задати наступні параметри:

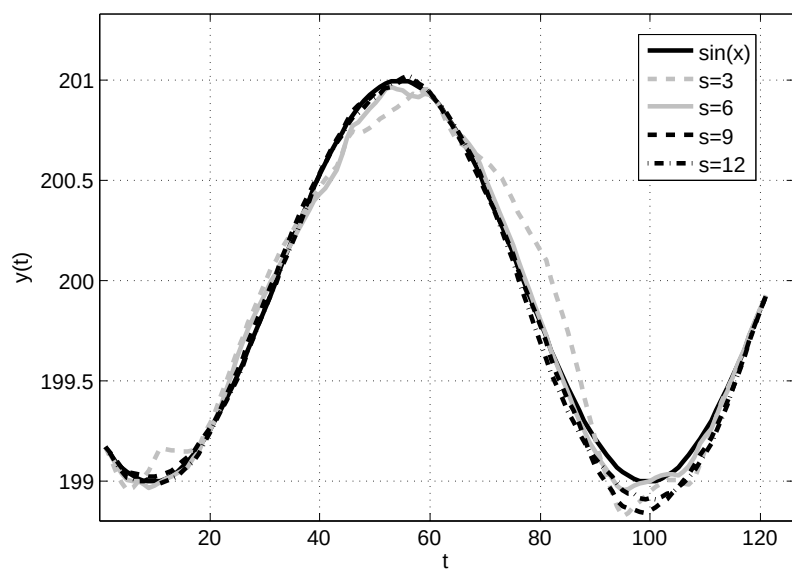


Рис. 2.2. Дискретизація та відновлення $y = \sin(x)$. Складна ієрархія

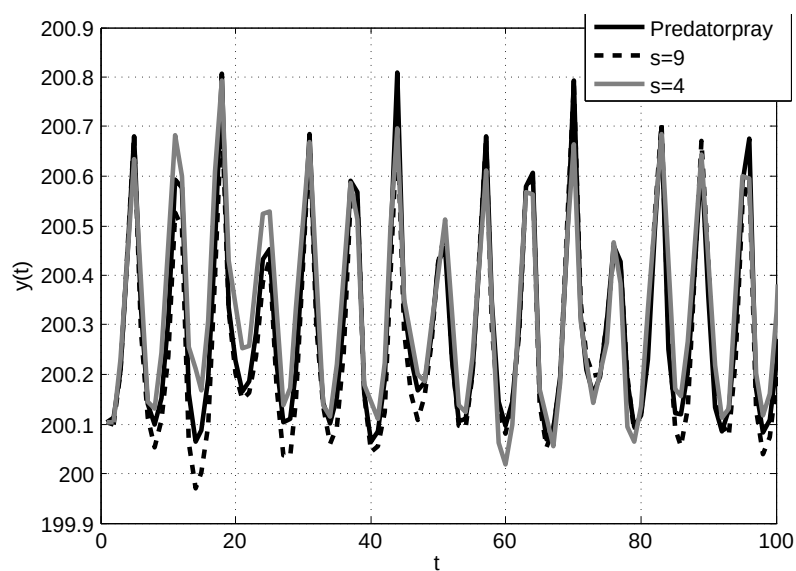


Рис. 2.3. Дискретизація та відновлення для моделі Хижак-жертва. Проста ієрархія

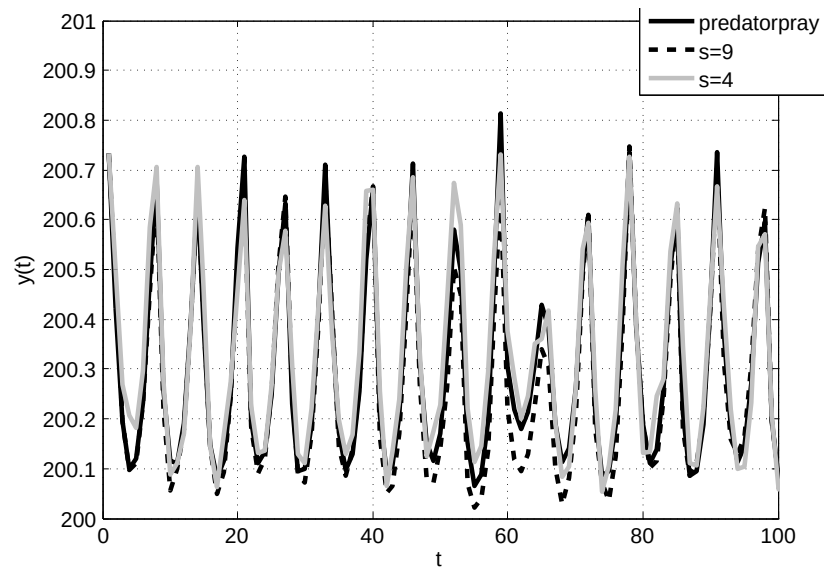


Рис. 2.4. Дискретизація та відновлення для моделі Хижак-жертва. Складна ієрархія

- 1) Вид ієрархії приростів часу (проста – степені двійки, складна – добуток степенів перших простих чисел).
- 2) Величини s – кількість станів та r – порядок ланцюга Маркова. Ці параметри можна зробити індивідуальними для кожного рівня дискретизації, знаходження оптимальних параметрів буде обговорюватись у підрозділі 2.3.
- 3) Величина порогу δ , та мінімальна кількість представників кожного стану СЛМ N_{min} .

Отже, алгоритм побудови прогнозу зводиться до наступних кроків:

1. Визначити ієрархію приростів часу – послідовність величин Δt , максимальне з яких повинне відповідати довжині навчальної вибірки N .
2. Для кожного приросту часу Δt у порядку зростання приростів, провести прогнозування станів та відновити за цими станами ряд.

При цьому треба:

- (а) обчислити прирости з дискретизацією Δt ;

- (b) перевести ряд приростів у ряд номерів станів $\{1, \dots, s\}$;
- (c) обчислити ймовірності переходів для узагальнених станів;
- (d) побудувати ряд прогнозних станів за допомогою застосування процедури визначення найбільш ймовірного наступного стану;
- (e) відновити ряд прогнозних значень з ряду подій з дискретизацією Δt ;
- (f) провести процедуру склеювання результату з рядом, який отримався в результаті склеювання попередніх шарів (з меншим кроком Δt). У випадку, якщо даний ряд є першим, повернути його без змін у якості результату склеювання.

3. Останній склеєний ряд склеїти з продовженням лінійного тренду, побудованого по всім попередньо відомим точкам.

Ряд, склеєний з лінійним трендом, є результатом прогнозування.

2.2. Складні ланцюги Маркова та вплив передісторії (пам'яті). Порядок ланцюга Маркова

У випадку простого ланцюга Маркова на основі ряду станів обчислюються ймовірності переходів з кожного стану у всі інші. Ці ймовірності прийнято записувати у вигляді квадратної матриці розмірністю s . Наприклад, зв'язавши рядки матриці зі станами, з яких відбувається перехід, а стовбці (елементи цих рядків) – зі станами, у які визначається ймовірності переходу.

Наприклад, для п'яти станів можна записати матрицю ймовірностей переходів A :

$$A = \begin{pmatrix} p_{11} & p_{12} & p_{13} & p_{14}p_{15} \\ p_{21} & p_{22} & p_{23} & p_{24}p_{25} \\ p_{31} & p_{32} & p_{33} & p_{34}p_{35} \\ p_{41} & p_{42} & p_{43} & p_{44}p_{45} \\ p_{51} & p_{52} & p_{53} & p_{54}p_{55} \end{pmatrix},$$

в якій, наприклад $A(1, 2) = p_{12}$ – ймовірність переходу із стану 1 в стан 2.

Якщо відомий вектор $p_t = (1, 0, 0, 0, 0)$, ймовірностей виникнення станів у поточний момент t то, згідно теореми Байеса, вектор ймовірностей наступного стану обчислюється за формулою:

$$p_{t+1} = p_t \cdot A,$$

а вектор ймовірностей станів через k кроків обчислюється за формулою:

$$p_{t+k} = p_t \cdot A^k,$$

У випадку складного ланцюга Маркова, ймовірність наступного стану залежить не тільки від попереднього стану, а й від послідовності r станів, які відбулись перед даним. У цьому випадку, необхідно обчислити ймовірності переходів з послідовності r станів у стан $(r + 1)$. Формально ці ймовірності можна було б записати в прямокутну таблицю розмірності (r^s, s) .

Можна звести ланцюги Маркова порядку r до ланцюга порядку 1, узагальнивши поняття “наявний стан”, включивши у нього послідовність з r станів, які передують стану, ймовірність появи якого обчислюється. Таким чином, ймовірності переходів можна записати в квадратну матрицю розмірністю (r^s, r^s) .

Наведемо приклад матриці ймовірностей переходів A для ланцюга Маркова порядку $r = 2$ та з кількістю станів $s = 3$. Відповідність

послідовностей станів введеним узагальненим станам наступна:

Послідовність станів	Узагальнений стан
11	1
12	2
13	3
21	4
22	5
23	6
31	7
32	8
33	9

Тоді матриця A ймовірностей переходів для узагальнених станів матиме розмірність 99, наприклад:

$$A = \begin{pmatrix} p_{11,11} & 0 & 0 & p_{21,11} & 0 & 0 & p_{31,11} & 0 & 0 \\ p_{11,12} & 0 & 0 & p_{21,12} & 0 & 0 & p_{31,12} & 0 & 0 \\ p_{11,13} & 0 & 0 & p_{21,13} & 0 & 0 & p_{31,13} & 0 & 0 \\ 0 & p_{12,21} & 0 & 0 & p_{22,21} & 0 & 0 & p_{32,21} & 0 \\ 0 & p_{12,22} & 0 & 0 & p_{22,22} & 0 & 0 & p_{32,22} & 0 \\ 0 & p_{12,23} & 0 & 0 & p_{22,23} & 0 & 0 & p_{32,23} & 0 \\ 0 & 0 & p_{13,31} & 0 & 0 & p_{23,31} & 0 & 0 & p_{33,31} \\ 0 & 0 & p_{13,32} & 0 & 0 & p_{23,32} & 0 & 0 & p_{33,32} \\ 0 & 0 & p_{13,33} & 0 & 0 & p_{23,33} & 0 & 0 & p_{33,33} \end{pmatrix}.$$

Процес прогнозування полягає в наступному: вибирається останній стан (у випадку ланцюга Маркова порядку $r > 1$ береться послідовність r останніх станів). Визначається ймовірність переходу з даного стану у всі можливі. Із усіх можливих наступних станів вибирається стан з максимальною ймовірністю. У випадку декількох станів з максимальною ймовірністю процес прийняття рішення буде описаний далі.

Вибраний найбільш ймовірний стан приймається, як наступний прогнозований стан та процедура повторюється з наступним (доданим як останній) станом. Таким чином, отримуємо ряд прогнозних станів для даної величини дискретизації Δt .

Далі по одержаному ряду станів та відомому початковому значенню y_N відбувається відновлення ряду для даної дискретизації Δt . При цьому кожний стан являє собою Δt точок ряду. З кожним станом зв'язаний середній приріст, який додається до значення останньої точки ряду та обчислюється наступна дискретизована точка. Проміжні точки заповнюються, як лінійна інтерполяція двох відомих сусідніх точок.

Розглянемо процедуру покрокового прогнозування, при якій у якості наступного прогнозного стану вибирається стан з найбільшою ймовірністю при даних умовах. Для цього ми використовуємо матрицю ймовірностей переходів із стану в стан. При цьому необхідно враховувати, що ймовірності можуть обчислюватись тільки наближено. Для підвищення точності оцінювання ймовірностей необхідно збільшити об'єм вибірки можливих переходів між станами. Але дана вибірка обмежена конкретною реалізацією ряду, що прогнозується, тому пропонується ряд алгоритмічних правил, які дозволяють врахувати принципову наближеність оцінених ймовірностей переходів при прогнозуванні майбутніх станів дискретного ланцюга Маркова.

Розглянемо конкретний випадок розподілу ймовірностей наступного стану та проілюструємо алгоритм вибору прогнозного стану при заданих умовах.

Як видно з рис.2.5, для того, щоб не допустити втрати станів, ймовірності яких обчислені з похибкою (наприклад, стан 6 на рис.2.5), необхідно до стану з максимальною ймовірністю додати сусідні стани, ймовірність яких не відрізняється від максимального більше ніж на δ .

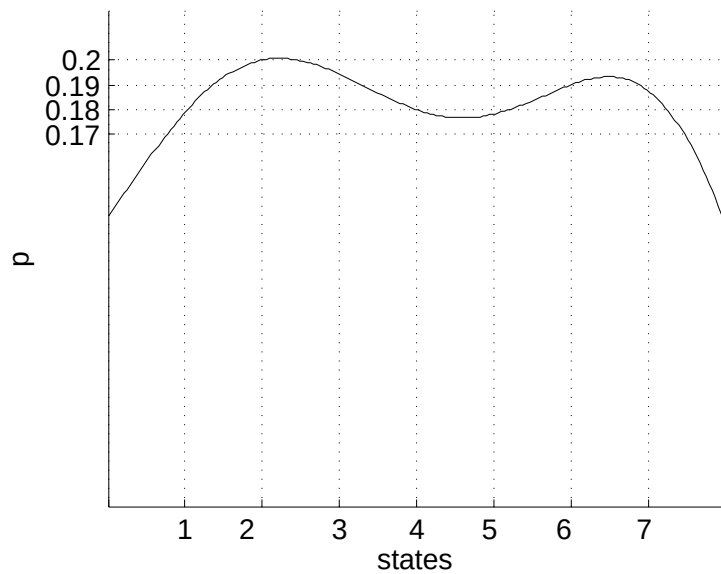


Рис. 2.5. Розподіл ймовірностей на певному кроці

Параметр δ вибирається залежно від похибки оцінювання ймовірностей, його значення необхідно експериментально апробувати.

Якщо прийняти $\delta = 0,01$ (рис.2.5), то до множини найбільш ймовірних станів увійдуть наступні: $\{2, 3, 6, 7\}$. Декілька сусідніх станів з максимальною ймовірністю будемо називати кластером. У даному випадку це стани 2 та 3. Якщо прийняти $\delta = 0,02$, тоді множина найбільш ймовірних станів буде $\{1, 2, 3, 4, 6, 7\}$. Для прогнозування необхідно вибрати один або два найбільш ймовірних стани. Вибираючи два стани з найбільш ймовірних, метод дозволяє розглянути два найбільш ймовірних сценарії можливого продовження ряду. Для вибору наступного стану пропонуються наступні правила:

1. Якщо рівні (квантовані прирости) з максимальною ймовірністю створюють декілька кластерів (кластер - група з декількох сусідніх рівнів – елементів кластера, мінімальний кластер – один ізольований рівень), то вибираємо найбільший за розміром кластер.

2. Якщо число елементів кластера непарне, то вибираємо центральний елемент у якості k_{max} .
3. Якщо число елементів в кластері парне, то розглядаємо два центральні елементи кластеру, та вибираємо в якості k_{max} той з них, який ближче до центру розподілу.
4. Якщо два центральних елементи кластера знаходяться на однаковій відстані від центру розподілу, то необхідно розглянути обидва, як можливі варіанти 1 та 2 значень k_{max} (точка біфуркації).
5. Якщо кластерів з максимальним розміром декілька, то розглядаємо їх, як нові елементи, які, в свою чергу, можуть формувати кластери, з якими необхідно вчинити аналогічно пунктам 1–4.

В цей принцип відбору закладені наступні міркування:

1. Якщо існують два однакових сусідні стани з максимальною ймовірністю, то краще взяти той, який ближче до центра розподілу, щоб звести до мінімуму ризик відображення у прогнозі хибних трендів.
2. Якщо рівні з максимальною ймовірністю не сусідні (багатомодальний розподіл), то необхідно розглянути як мінімум два варіанти, оскільки це може бути зв'язане з біфуркаціями, які теж хотілось не упустити з розгляду.
3. Реалізуючи прогноз по варіанту 1) (на усіх етапах ієрархії), отримується в певному наближенні нижня границя прогнозу, а реалізуючи прогноз за варіантом 2) – його верхня границя.

У прикладі розподілу (рис. 2.5) у випадку $\delta = 0,02$ кластерами найбільш ймовірних станів будуть множини $\{1, 2, 3, 4\}$ та $\{6, 7\}$. Згідно правилам 1 та 3, рекомендується вибрати середній стан найбільшого (у нашому випадку, першого) кластеру - стан 3 (згідно правилу 4, із двох середніх 2 та 3 найближчий до центру розподілу - стан 4).

Якщо прийняти ($\delta = 0,01$), кластерами найбільш ймовірних станів

будуть множини $\{2, 3\}$ та $\{6, 7\}$. Проміжок між станами 4 та 5 необхідно взяти як центр розподілу (стан, що відповідає нульовому відхиленню, мав би знаходитись між станами 4 та 5), таким чином відстані від центра розподілу до двох вибраних кластерів є однаковими. В такому випадку має місце бімодальний розподіл (обидва кластери є однаково ймовірними), тому для адекватного моделювання майбутніх сценаріїв, необхідно розглянути 2 варіанти. Перший сценарій буде починатись зі стану 3, (згідно правила 3 як найближчий до центра розподілу у кластері з парної кількості станів) а другий - зі стану 6 (за аналогічним правилом 3).

Отже, у випадку невизначеності майбутнього стану, обмежуємось розглядом одного чи двох варіантів (сценаріїв) прогнозу: один з варіантів є верхньою границею прогнозу, а другий – нижньою.

2.3. Ієрархія приростів часу та процедура склеювання

По одержаному ряду станів та відомому початковому значенню відбувається відновлення ряду для даної дискретизації Δt . При цьому кожний стан являє собою Δt точок ряду. На етапі класифікації станів з кожним станом був зв'язаний середній приріст $r_{avg,i}$, який додається до значення останньої точки ряду та обчислюється наступна дискретизована точка. Проміжні точки заповнюються як лінійна інтерполяція відомих двох сусідніх точок. Алгоритм відновлення значень ряду y_t на основі початкової ціни p_t та ряду середніх приростів, $r_{avg,ik}$ відповідних прогнозним станам

s_k можна задати послідовністю обчислень (6):

$$\begin{aligned}
 y_t &= p_t \\
 y_{t+1} &= y_t + r_{avg,i1}/\Delta t = p_t + r_{avg,i1}/\Delta t \\
 y_{t+2} &= y_{t+1} + r_{avg,i1}/\Delta t = p_t + r_{avg,i1}/\Delta t \\
 &\dots \\
 y_{t+\Delta t-1} &= y_{t+\Delta t-2} + r_{avg,i1}/\Delta t = p_t + (\Delta t - 1)r_{avg,i1}/\Delta t \\
 y_{t+\Delta t} &= y_{t+\Delta t-1} + r_{avg,i1}/\Delta t = p_t + \Delta t r_{avg,i1}/\Delta t = p_t + r_{avg,i1} \\
 y_{t+\Delta t+1} &= y_{t+\Delta t} + r_{avg,i2}/\Delta t = p_t + r_{avg,i1} + r_{avg,i2}/\Delta t \\
 &\dots
 \end{aligned}$$

$$y_{t+n\Delta t} = y_{t+\Delta t-1} + r_{avg,in}/\Delta t = p_t + \sum_{k=1}^n r_{avg,ik}, \quad (2.5)$$

де y_t - прогнозний ряд, $r_{avg,ik}$ - середній приріст, що відповідає i -му спрогнозованому дискретному стану. Дана процедура є аналогом чисельного розв'язання задачі Коші. Відомо, що метод Ейлера [?] для задачі Коші суттєвим недоліком має накопичення похибки з кожною наступною ітерацією. Для подолання цієї проблеми були запропоновані багатокрокові методи (Рунге-Кутти), які використовують декілька попередніх точок для обчислення кожної наступної. У даній дисертації пропонується альтернативний підхід до вирішення цієї проблеми, ідея якого полягає у коригуванні накопиченої похибки за допомогою прогнозу, виконаного з більш крупним кроком. Наприклад, прогноз значення $y_{t+2\Delta t}$ можна отримати як два кроки прогнозування Δt , а також, застосувавши цей же алгоритм до вихідного ряду з кроком дискретизації $2\Delta t$ та здійснивши одну ітерацію прогнозування. У першому випадку похибка буде накопичена за два кроки, а у другому – тільки за один. Це дає можливість підкоректувати прогнозну точку $y_{t+2\Delta t}$, отриману від двокрокового прогнозування так, щоб її значення співпало з більш точним прогнозом цієї точки, зробленим з кроком $2\Delta t$, а усі інтерпольовані точки були

підкоректовані відповідно. Таку процедуру корекції прогнозного часового ряду назвемо склеюванням. Дана процедура організована ітераційно та проводиться, починаючи з менших приростів, додаючи на кожному кроці прогноз з більшим приростом часу.

Прогнози часового ряду необхідно обчислювати з різними кроками дискретизації часу. Наприклад, аналогічно з дискретним перетворенням Фур'є, розглядаємо прирости величиною степенів двійки. Спочатку обчислюємо прирости, як різницю двох сусідніх значень ряду, потім здійснюємо розрідження ряду з кроками 2, 4, 8, 16 і т. д. Позначимо поточну різницю між вимірами у часі через Δt .

Для кожного Δt виконуємо перетворення ряду приростів у ряд станів, здійснюємо прогнозування майбутньої послідовності станів, потім відновлюємо ряд із заданою дискретизацією по спрогнозованому ряді станів.

Ряди, отримані при відновленні для різних Δt , проходять ітераційну процедуру склеювання, в результаті якої прогнози з більшим кроком Δt коригують значення ряду, отриманих при прогнозуванні з меншими значеннями інтервалу Δt .

Таким чином, вибирається ієрархія приростів, кожен з яких відповідає за свою частоту, на якій відбувається прогнозування та поєднання прогнозів на різних Δt в один ряд у процесі склеювання.

Суть процесу склеювання полягає в наступному. Процедура склеювання є ітераційною, ряд з кожною наступною (з більшим кроком) дискретизацією коригує, підтягуючи до своїх вузлових точок прогноз, отриманий після склеювання на попередніх, менших Δt . Перетворення, які виконуються в процесі склеювання, можна записати у вигляді наступних обчислень:

Нехай задана ієрархія відносних (квантованих) приростів:

$y_{\Delta t_k}^k, y_{2\Delta t_k}^k, \dots, y_N^k; \quad k = 1, 2, \dots, K_{max}$ для кроків дискретизації $\Delta t_1 = 1; \Delta t_2, \dots, \Delta t_{K_{max}}$ відповідно.

Нехай процедура склеювання проведена для сукупності кроків $\Delta t_1 = 1; \Delta t_2; \dots; \Delta t_{K_{max}}$ та отриманий ряд значень $X_0^k, X_1^k, \dots, X_N^k$ з інтервалом дискретизації Δt_k . Тоді продовження процедури склеювання та перехід до ряду значень $X_0^{k-1}, X_1^{k-1}, \dots, X_N^{k-1}$ відбувається за алгоритмом:

початкове значення:

$$x_0^{k+1} = x_0^k = x_0;$$

для проміжку $0 < t \leq \Delta t_{k+1}$:

$$\begin{aligned} x_1^{k+1} &= x_0^{k+1} + (x_1^k - x_0^k) + 1 \cdot \frac{x_0^{k+1} y_{\Delta t_{k+1}}^{k+1} - (x_{\Delta t_{k+1}}^k - x_0^k)}{\Delta t_{k+1}}; \\ x_2^{k+1} &= x_0^{k+1} + (x_2^k - x_0^k) + 2 \cdot \frac{x_0^{k+1} y_{\Delta t_{k+1}}^{k+1} - (x_{\Delta t_{k+1}}^k - x_0^k)}{\Delta t_{k+1}}; \\ &\dots \\ x_{\Delta t_{k+1}-1}^{k+1} &= x_0^{k+1} + (x_{\Delta t_{k+1}-1}^k - x_0^k) + \\ &+ (\Delta t_{k+1} - 1) \cdot \frac{x_0^{k+1} y_{\Delta t_{k+1}}^{k+1} - (x_{\Delta t_{k+1}}^k - x_0^k)}{\Delta t_{k+1}}; \\ x_{\Delta t_{k+1}}^{k+1} &= x_0^{k+1} + (x_{\Delta t_{k+1}}^k - x_0^k) + \Delta t_{k+1} \cdot \frac{x_0^{k+1} y_{\Delta t_{k+1}}^{k+1} - (x_{\Delta t_{k+1}}^k - x_0^k)}{\Delta t_{k+1}} \equiv \\ &\equiv x_0^{k+1} (1 + y_{\Delta t_{k+1}}^{k+1}); \end{aligned}$$

для проміжку $\Delta t_{k+1} < t \leq 2\Delta t_{k+1}$:

$$\begin{aligned}
x_{\Delta t_{k+1}+1}^{k+1} &= x_{\Delta t_{k+1}}^{k+1} + \left(x_{\Delta t_{k+1}+1}^k - x_{\Delta t_{k+1}}^k \right) + \\
&+ 1 \cdot \frac{x_{\Delta t_{k+1}}^{k+1} y_{2\Delta t_{k+1}}^{k+1} - (x_{2\cdot\Delta t_{k+1}}^k - x_{\Delta t_{k+1}}^k)}{\Delta t_{k+1}}; \\
x_{\Delta t_{k+1}+2}^{k+1} &= x_{\Delta t_{k+1}}^{k+1} + \left(x_{\Delta t_{k+1}+2}^k - x_{\Delta t_{k+1}}^k \right) + \\
&+ 2 \cdot \frac{x_{\Delta t_{k+1}}^{k+1} y_{2\Delta t_{k+1}}^{k+1} - (x_{2\cdot\Delta t_{k+1}}^k - x_{\Delta t_{k+1}}^k)}{\Delta t_{k+1}}; \\
&\dots \\
x_{2\Delta t_{k+1}-1}^{k+1} &= x_{\Delta t_{k+1}}^{k+1} + \left(x_{2\Delta t_{k+1}-1}^k - x_{\Delta t_{k+1}}^k \right) + \dots \\
&+ (\Delta t_{k+1} - 1) \cdot \frac{x_{\Delta t_{k+1}}^{k+1} y_{2\Delta t_{k+1}}^{k+1} - (x_{2\cdot\Delta t_{k+1}}^k - x_{\Delta t_{k+1}}^k)}{\Delta t_{k+1}}; \\
x_{2\Delta t_{k+1}}^{k+1} &= x_{\Delta t_{k+1}}^{k+1} + \left(x_{2\Delta t_{k+1}}^k - x_{\Delta t_{k+1}}^k \right) + \\
&+ \Delta t_{k+1} \cdot \frac{x_{\Delta t_{k+1}}^{k+1} y_{2\Delta t_{k+1}}^{k+1} - (x_{2\cdot\Delta t_{k+1}}^k - x_{\Delta t_{k+1}}^k)}{\Delta t_{k+1}} \equiv \\
&\equiv x_{\Delta t_{k+1}}^{k+1} \left(1 + y_{2\Delta t_{k+1}}^{k+1} \right);
\end{aligned}$$

для останнього проміжку $N - \Delta t_{k+1} < t \leq N$:

$$\begin{aligned}
x_{N-\Delta t_{k+1}+1}^{k+1} &= x_{N-\Delta t_{k+1}}^{k+1} + \left(x_{N-\Delta t_{k+1}+1}^k - x_{N-\Delta t_{k+1}}^k \right) + \\
&+ 1 \cdot \frac{x_{N-\Delta t_{k+1}}^{k+1} y_N^{k+1} - (x_N^k - x_{N-\Delta t_{k+1}}^k)}{\Delta t_{k+1}}; \\
x_{N-\Delta t_{k+1}+2}^{k+1} &= x_{N-\Delta t_{k+1}}^{k+1} + \left(x_{N-\Delta t_{k+1}+2}^k - x_{N-\Delta t_{k+1}}^k \right) + \\
&+ 2 \cdot \frac{x_{N-\Delta t_{k+1}}^{k+1} y_N^{k+1} - (x_N^k - x_{N-\Delta t_{k+1}}^k)}{\Delta t_{k+1}}; \\
&\dots \\
x_{N-1}^{k+1} &= x_{N-\Delta t_{k+1}}^{k+1} + \left(x_{N-1}^k - x_{N-\Delta t_{k+1}}^k \right) + \\
&+ (\Delta t_{k+1} - 1) \cdot \frac{x_{N-\Delta t_{k+1}}^{k+1} y_N^{k+1} - (x_N^k - x_{N-\Delta t_{k+1}}^k)}{\Delta t_{k+1}}; \\
x_N^{k+1} &= x_{N-\Delta t_{k+1}}^{k+1} + \left(x_N^k - x_{N-\Delta t_{k+1}}^k \right) + \\
&+ \Delta t_{k+1} \cdot \frac{x_{N-\Delta t_{k+1}}^{k+1} y_N^{k+1} - (x_N^k - x_{N-\Delta t_{k+1}}^k)}{\Delta t_{k+1}} \equiv \\
&\equiv x_{N-\Delta t_{k+1}}^{k+1} (1 + y_N^{k+1}).
\end{aligned}$$

Цю процедуру проводимо для $k = 1, 2, 3, \dots, K_{\max}$, починаючи з $k = 1$. Початковий ряд (при $k = 1$ для $\Delta t_1 = 1$) $x_0^1, x_1^1, \dots, x_N^1$ знаходимо за алгоритмом:

$$\begin{aligned}
x_0^1 &= x_0; \\
x_1^1 &= x_0^1 (1 + y_1^1); \\
x_2^1 &= x_1^1 (1 + y_2^1); \\
&\dots \\
x_N^1 &= x_{N-1}^1 (1 + y_N^1).
\end{aligned}$$

Розглянемо процедуру склеювання часового ряду на прикладі фрагменту індексу PFTS:

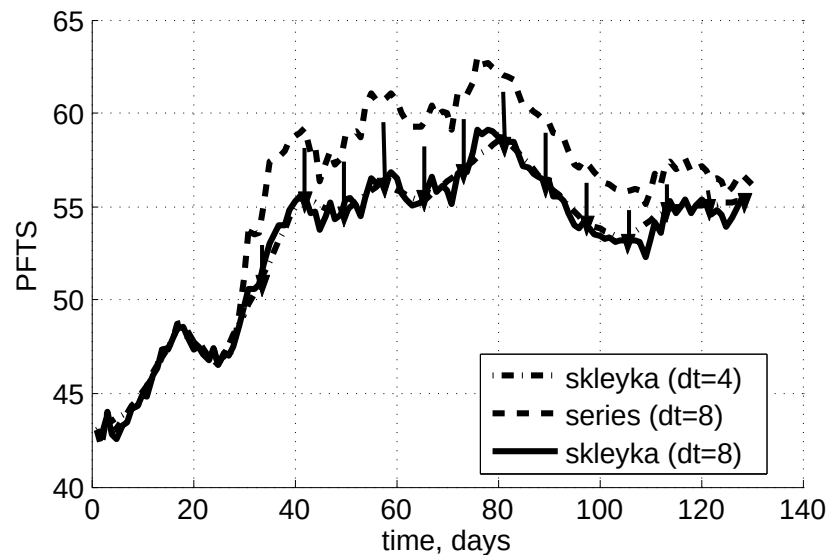


Рис. 2.6. Процес склеювання двох ієрархій часового ряду

На рис.2.6 зображено процес склеювання прогнозів з інтервалами дискретизації $\Delta t = 1, 2, 4$ до нового ряду з дискретизацією $\Delta t = 8$. Як видно з рисунка, точки ряду, склеєного з $\Delta t=1, 2, 4$ (рис.2.6, пунктирна лінія) “підтягуються” до точок нового ряду з дискретизацією $\Delta t=8$ (рис.2.6, штрихпунктирна лінія). В результаті отримується ряд, склеєний з $\Delta t 1, 2, 4, 8$ (рис. 2.6, суцільна лінія). Як видно з рисунку, характерні риси прогнозу при $\Delta t = 4$ зберігаються після склеювання з прогнозним рядом при $\Delta t = 8$.

Ієрархія приростів дискретизацій, як степені числа 2 – не єдиний спосіб побудови ієрархії дискретизацій. Для покращення прогнозу можна використовувати більш часту ієрархію дискретизацій, наприклад, добуток степенів простих чисел.

Висувається гіпотеза, що при більшій кількості прогнозів з різними проміжками дискретизації Δt , отриманий після склеювання прогноз буде більш точно відображати характер процесу, що прогнозується. Тому у якості альтернативної до множини значень $\Delta t_i = 2^i$, пропонується множина $\Delta t_i = p_1^{m_1} \cdot p_2^{m_2} \cdot \dots \cdot p_n^{m_n}$, де $\{p_k\}$ - підмножина множини простих чисел, $\{m_k\}$ - спеціально вибрана множина показників степенів. Множина простих чисел $\{p_k\}$ та відповідна множина їх степенів $\{m_k\}$ вибираються з розрахунку, щоб кількість різних значень Δt_i було максимальним. У цьому випадку, згідно припущення, можна покрити фрагмент часового ряду, взятий як навчальна вибірка, більш густою сіткою приростів Δt_i , таким чином використавши більше інформації, наявної у часовому ряді.

Припустимо, що довжина ряду, який необхідно розкласти на ієрархії – N . Взявши послідовність перших простих чисел $p_1, p_2, \dots, p_n$ та послідовність їх степенів, при яких $p_1^{m_1} \cdot p_2^{m_2} \cdot \dots \cdot p_n^{m_n} = N_1 < N$. Число N_1 повинне бути якомога ближчим до N , щоб збільшити кількість точок, які використовуються при прогнозуванні. Прості числа і їх степені слід вибирати таким чином, щоб кількість дільників цього числа було максимальним. Це забезпечить максимальну щільність сітки точок дискретизацій. Оскільки N_1 є добутком заданих степенів простих чисел, то кількість дільників числа N_1 можна обчислити за наступною формулою: $(m_1 + 1) \cdot (m_2 + 1) \cdot \dots \cdot (m_n + 1)$. Ця кількість приростів Δt_i повинна бути максимізована.

Задача полягає в тому, щоб знайти послідовність простих чисел та відповідні їх степені, для яких, по-перше, добуток $p_1^{m_1} \cdot p_2^{m_2} \cdot \dots \cdot p_n^{m_n} = N_1$ був би близьким до N , а по-друге, кількість дільників цього числа була б

максимальною. Для розв'язання цієї задачі можна використовувати повний перебір варіантів, оскільки цей процес не займає значного обчислювального часу, а також проводиться лише один раз на початку прогнозування.

Розглянемо приклад:

Для $N = 5000$ необхідно підібрати такий набір простих чисел $2, 3, 5, 7, \dots, q_k$, при якому число кроків дискретизації буде близьким до максимально можливого. Максимальне k та відповідне йому q_k визначається умовою:

$$\frac{\ln N}{k \ln q_k} \approx 1; \quad \Rightarrow \quad k \approx \frac{\ln N}{\ln q_k}.$$

$$N = 5000; \quad q_1 = 2; \quad q_2 = 3; \quad q_3 = 5; \quad q_4 = 7; \quad q_5 = 11;$$

$$k = 3; \quad \frac{\ln 5000}{\ln 5} = 5,29; \quad k = 4; \quad \frac{\ln 5000}{\ln 7} = 4,37;$$

$$k = 5; \quad \frac{\ln 5000}{\ln 11} = 3,55.$$

Вибираємо $k = 4$:

$$m_2 = \frac{\ln 5000}{4 \ln 2} = 3,07; \quad m_3 = \frac{\ln 5000}{4 \ln 3} = 1,94; \quad ;$$

$$m_7 = \frac{\ln 5000}{4 \ln 7} = 1,09; \quad \Rightarrow$$

$$2^3 3^2 5^1 7^1 = 2520; \quad \Rightarrow \quad N = 2^4 3^2 5^1 7^1 = 5040; \quad \Rightarrow$$

$$K_{\max} = (m_2 + 1)(m_3 + 1)(m_5 + 1)(m_7 + 1) = 5 \cdot 3 \cdot 2 \cdot 2 = 60; \quad \Rightarrow$$

$$\Delta t = 1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 12, 14, 15, 16, 18, 20, 21, 24, 28, 30,$$

$$35, 36, 40, 42, 45, 48, 56, 60, 63, 70, 72, 80, 84, 90, 112, 120, 126,$$

$$140, 144, 168, 180, 210, 240, 252, 280, 315, 336, 360, 420, 504, 560,$$

$$630, 720, 840, 1008, 1260, 1680, 2520, 5040.$$

Таким чином, замість 12 приростів у випадку простої ієрархії ($2^{12}=4096$), маємо 60 приростів, що дозволить створити більш густу сітку часових приростів, та за висунутим припущенням, отримати більш точний прогноз.

Вибір оптимальних параметрів алгоритму прогнозування

Описаний вище алгоритм має параметри, які необхідно підбирати експериментально. Для цього проводиться процедура оптимізації параметрів, проводячи тестові прогнози для відомої частини ряду та оцінюючи результат.

Вихідний ряд розбивається на 2 частини, першу приймаємо як навчальну вибірку, а іншу як тестову. Будемо проводити тестові прогнози при різних значеннях параметрів та оцінювати відхилення прогнозу від реального значення.

Оскільки висунуто припущення, що для кожного приросту часу необхідно окремо підбирати параметри r та s , то пропонується наступний алгоритм оптимізації параметрів:

1. Будуємо перше наближення – лінійний тренд.
2. Для кожного приросту часу Δt (дискретизації) в порядку спадання, проводимо підбір оптимальних параметрів. Для цього:
 - (a) Для кожної пари параметрів (r, s) , при яких будуть перевірятись прогнози, та для декількох початкових точок виконуємо наступні дії:
 - i. Обчислюємо прогноз на даному частотному рівні та здійснюємо склеювання для усіх раніше побудованих прогнозів з більшим кроком Δt , використовуючи вже підібрані на минулих кроках оптимальні r та s .
 - ii. Порівнюємо прогноз на величину Δt з реальними значеннями (тестовий проміжок). Обчислюємо похибку, як суму квадратів відхилень відповідних точок прогнозу та реальних точок.
 - (b) Таку похибку обчислюємо для декількох прогнозів з різних точок. Знаходимо середнє значення похибок для даної пари r та s .

- (с) Із усіх пар r та s вибираємо ту, яка в середньому дає найменшу похибку.
- (d) Приймаємо кращу пару для даного кроку приросту часу Δt та переходимо до наступного рівню приросту в порядку спадання.

Таким чином, для кожного рівня приросту часу знаходимо оптимальні параметри, які будуть використовуватись в остаточному прогнозі.

Припускається, що для певних рядів існує закономірність в оптимальних параметрах. Для виявлення цієї закономірності проводиться експериментальна робота, мета якої - отримати оптимальний набір параметрів для ряду без повного перебору варіантів.

2.4. Прогнозування часових рядів з точним періодом

В даному параграфі розглядаються результати апробації запропонованого методу на рядах з точним періодом (синусоїди з різними частотами) та рядах детермінованого хаосу (ряди динаміки хижаків та жертв в моделі Вольтерра-Лотки). При експериментах на рядах з точним періодом, наприклад, залежності $y(t)=\sin(at+b)$, оптимальний порядок ланцюга Маркова r прямує до максимально можливого значення, оскільки ряд станів містить абсолютно періодичні фрагменти. Але цей параметр обмежений тим, що для визначення значення елементів матриці ймовірностей переходів кожен перехід між станами повинен бути представлений достатньою кількістю повторень у ряді станів. Для цього необхідно, щоб кількість елементів матриці була б однакового порядку з розміром вибірки станів. Якщо припустити, що в середньому кількість переходів між станами на кожен комірку матриці повинні бути k а розмір вибірки станів дорівнює N , тоді оптимальне співвідношення для параметрів r та s можна записати рівнянням [49] :

$$k \cdot r^{s+1} \approx N. \quad (2.6)$$

Якщо прийняти факт, що для рядів з точним періодом порядок ланцюга Маркова r повинен бути максимальним, тоді можна здійснити перебір тільки для параметрів s та k . Експерименти показують, що параметр k для синусоїд не повинен перевищувати кількості коливань, представлених у вибірці навчання.

Розглянемо прогнози синусоїди з різним періодом.

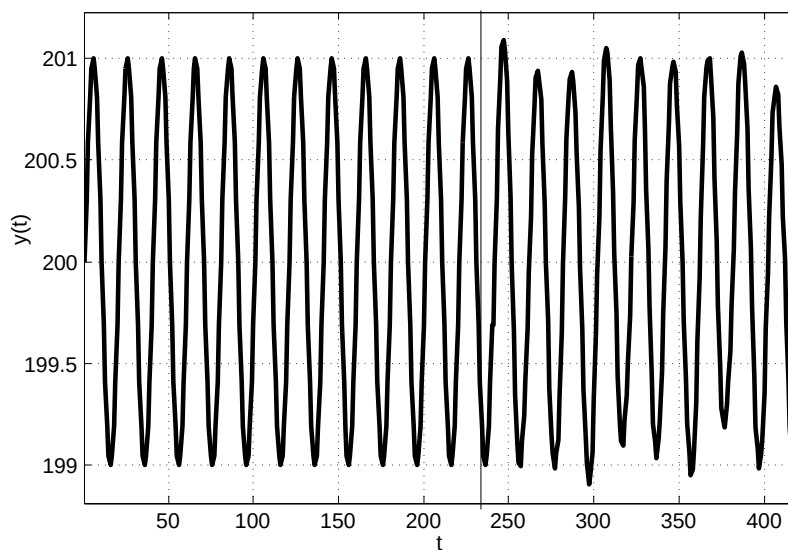


Рис. 2.7. Прогнозування $y = \sin\left(\frac{2\pi}{20}t\right)$ при $k = 5$, $s = 9$

Неточності прогнозування, як видно з рис. 2.8, обумовлені похибкою квантування на Δt , більших періоду синусоїди. На рис. 2.7 у вибірці навчання представлена достатня кількість періодів для відновлення на великих ієрархіях, а на рис. 2.8 цих представників недостатньо. Не зважаючи на це, перше коливання синусоїди з періодом 60 точок було спрогнозоване. З цього можна зробити висновок, що максимальний горизонт прогнозу може бути не більше 10-15 % від довжини навчальної вибірки.

Прогнозування рядів динамічного хаосу

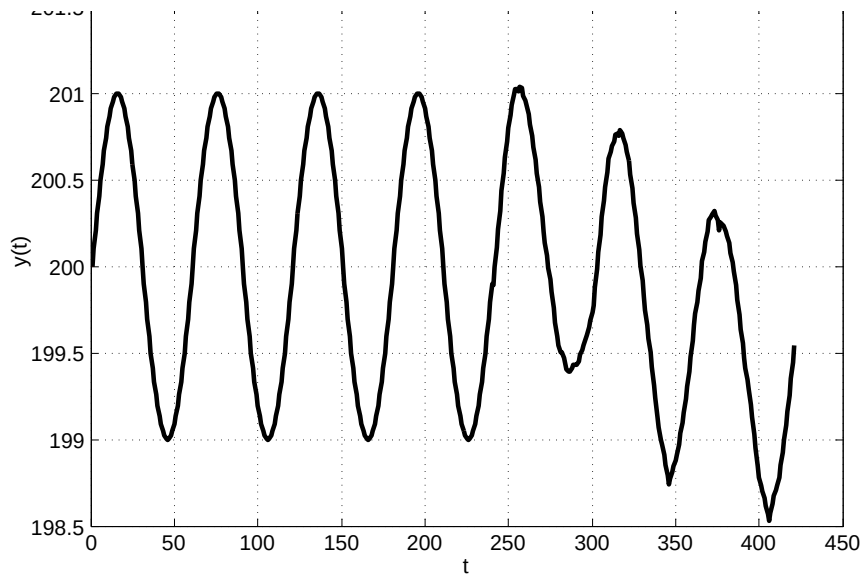


Рис. 2.8. Прогнозування $y = \sin\left(\frac{2\pi}{60}t\right)$ при $k = 2$, $s = 9$

Для налаштування методики, за допомогою моделі детермінованого хаосу типу “Хижак-жертва” з параметрами, близькими до критичного режиму, були згенеровані тестові ряди. Дані ряди мають коротку пам’ять, так як генеруються рекурентними співвідношеннями. На рис.2.9 рис.2.10 зображено прогноз таких рядів.

На рис.2.10 зображено прогноз при динамічному зростанні s пропорційно зростанню величини Δt . Причиною такого експерименту стало те, що для великих приростів часу необхідна велика точність, при похибках на такому рівні виникають биття, які помітні на обох графіках. Очевидно, що в другому випадку вдалося покращити прогноз для перших прогнозних значень.

Експериментальна робота з рядами динамічного хаосу показала, що запропонована методика дозволяє виявляти закономірності у модельних часових рядів, навіть якщо такі закономірності приховані від статистичних методів (прогнозування рядів детермінованого хаосу).

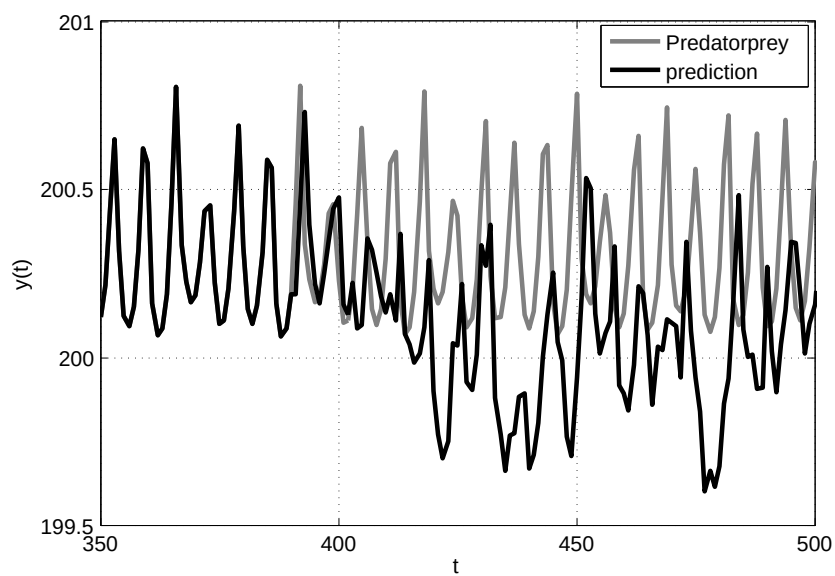


Рис. 2.9. Прогнозування ряду моделі “Хижак-жертва” $s=5$

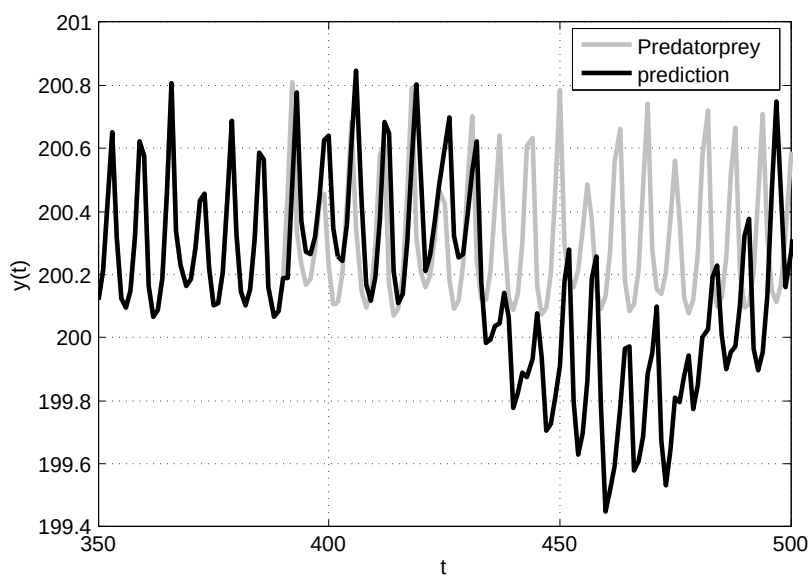


Рис. 2.10. Прогнозування ряду моделі “Хижак-жертва” $s = \sqrt{\Delta t}$

Виведене оптимальне співвідношення порядку ланцюгу Маркова r та кількості станів s , зроблено попередній підбір оптимальних параметрів.

2.5. Дискретне Фур'є-продовження низькочастотної складової рядів економічної динаміки.

Основна ідея даного розділу – дослідження частотних характеристик коливань, наявних у динамічному ряді, побудова моделі, яка апроксимує найбільш виражені коливання та екстраполяція побудованої моделі у майбутнє.

Задача прогнозування часових рядів на основі частотних складових розглядалась у ряді робіт [36, 101–103]. Зокрема, у роботі [36, 103] автор запропонував використання дискретного перетворення Фур'є для прогнозування врожайності. Перетворення Фур'є дає можливість дослідити амплітудно-частотну характеристику сигналу, що досліджується, але має певні недоліки, пов'язані з задачами прогнозування. Спектр частот, кратних двом, не дає можливості порівняно невеликою кількістю доданків ряду Фур'є представити сигнал з частотою, не кратною 2 до кроку часової дискретизації ряду. Підхід вейвлет-аналізу, запропонований у роботі [102], дає можливість подолати першу з вищезазначених проблем, але залишається проблема неефективності, збитковості кодування частотних складових. Нами запропоновано підхід, який дозволяє подолати зазначені проблеми, описуючи сигнал достатньо малою кількістю гармонік (до 10) зі змінною фазою та здійснити прогноз низькочастотної складової ряду за допомогою екстраполяції аналітичного представлення його частотних характеристик.

Апробація запропонованої у попередніх розділах методики показала, що прогнозування динамічних рядів за методикою складних ланцюгів Маркова показує достатньо високу точність для рядів із

високою регулярністю, повторюваністю значень для короткострокових прогнозів. При довгостроковому прогнозуванні також слід враховувати низькочастотні коливання у вихідному ряді. Для подолання цієї проблеми, у якості процедури попередньої підготовки ряду, пропонується апроксимація вихідного ряду сумою декількох низькочастотних гармонічних функцій. Далі проводиться детрендування ряду, тобто обчислення залишків апроксимації та застосування методу складних ланцюгів Маркова до залишків апроксимації. Прогноз залишків піддається оберненій процедурі відновлення ряду з використанням побудованого прогнозу залишків та екстраполяції гармонічної моделі тренду.

Метод побудови гармонічної моделі тренду може бути використаний як окремий метод прогнозування, його суть та результати застосування опубліковані у роботах [6, 12–14, 16].

Нехай ряд заданий послідовністю дискретних рівнів зі сталим кроком дискретизації часу Δt . Необхідно побудувати функцію $y(t)$, яка апроксимує задані емпіричні значення, оцінити параметри даної функції та експериментально перевірити ефективність прогнозів часового ряду на основі екстраполяції побудованої функції.

Пропонується низькочастотна апроксимуюча функція виду:

$$y_{abs}(t) = a + bt + \sum_{i=1}^m c_i \sin(d_i t + e_i), \quad (2.7)$$

або для відносного масштабу:

$$y_{rel}(t) = ae^{bt} \prod_{i=1}^m c_i \sin(d_i t + e_i). \quad (2.8)$$

Параметри моделі $a, b, c_1, c_2, \dots, c_m, d_1, \dots, d_m, e_1, \dots, e_m$ мають мінімізувати наступний критерій оптимальності:

$$F(a, b, c_1, \dots, c_m, d_1, \dots, d_m, e_1, \dots, e_m) = \sum_{i=1}^n (y_i - y_{lin}(t_0 + i\Delta t))^2, \quad (2.9)$$

або для відносного варіанту:

$$F(a, b, c_1, \dots, c_m, d_1, \dots, d_m, e_1, \dots, e_m) = \sum_{i=1}^n \left(1 - \frac{y_i}{y_{lin}(t_0 + i\Delta t)}\right)^2. \quad (2.10)$$

При розв'язуванні задачі оптимізації, необхідно задати початкові оцінки параметрів, які оптимізуються, а також накласти обмеження на їх значення. Суттєвим обмеженням необхідно піддати частоту d_i . У випадку короткої навчальної вибірки можлива ситуація, коли найкраща апроксимація відповідає невеликій частині повного гармонічного коливання. При цьому продовження даної аналітичної кривої може бути не характерним для ряду, що досліджується. Це проілюстровано на рис.2.11 та 2.12. Тому необхідною має бути умова відповідності не менш ніж половині повного коливання протягом навчальної вибірки, тобто частота коливання не може бути нижчою величини $\frac{\pi}{N}$, де N – довжина вибірки навчання. Також обмеження має бути накладене на високі частоти, що пояснюється, по-перше, складністю оптимізаційної задачі для високих частот, пов'язаною з наявністю великої кількості хибних локальних мінімумів, а по-друге, можливою некоректністю такого наближення на високих частотах. Емпірично було вибране обмеження 10 коливань за час навчальної вибірки.

Оскільки число параметрів функції F зростає зі збільшенням числа гармонік m , пропонується ітераційна апроксимація по одній (або двох) гармоніках, обчислення залишків та застосування наступної ітерації до більш високочастотних залишків для одночастинного наближення:

$$r_i(t) = r_{i-1}(t) - (a + bt + c_i \sin(d_i t + e_i)), \quad (2.11)$$

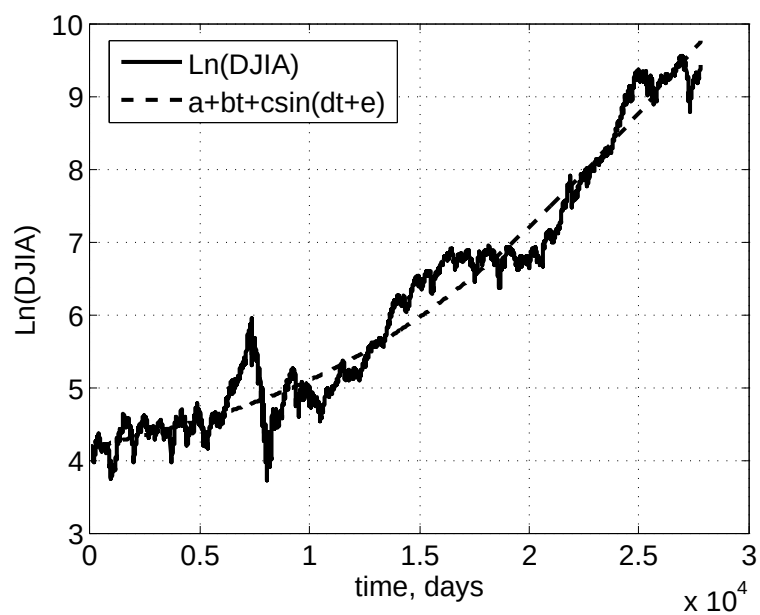


Рис. 2.11. Низькочастотна апроксимація індексу Dow Jones Industrial

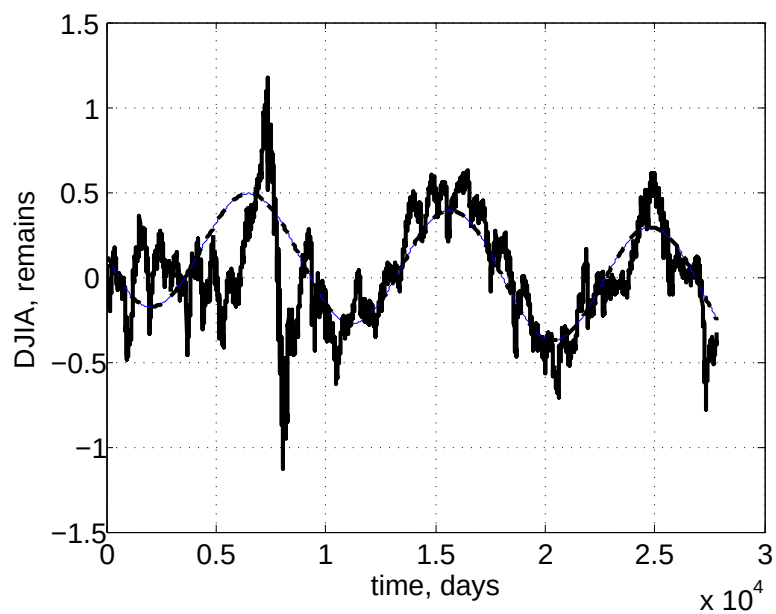


Рис. 2.12. Залишки апроксимації індексу DJIA та їх апроксимація

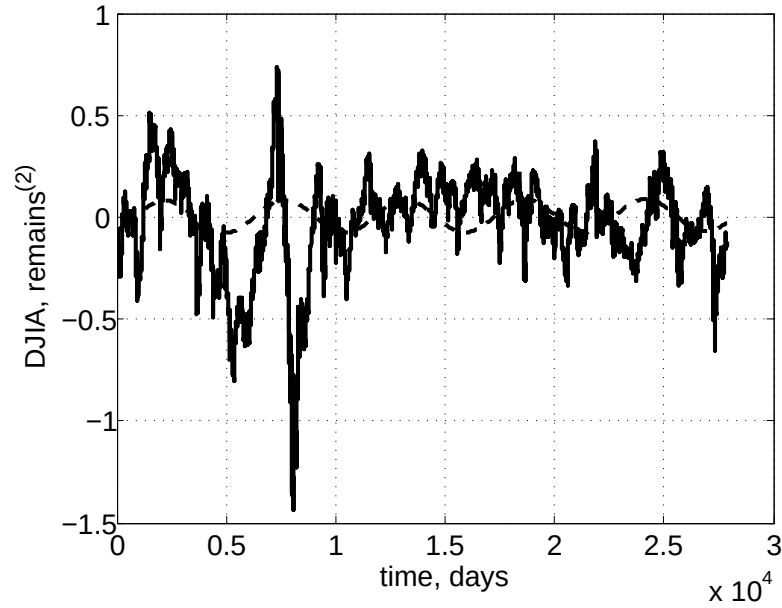


Рис. 2.13. Залишки другого порядку індексу DJIA та їх апроксимація або для m -частинного наближення:

$$r_i(t) = r_{i-1}(t) - \left(a + bt + \sum_{j=1}^m c_{ij} \sin(d_{ij}t + e_{ij}) \right). \quad (2.12)$$

Мінімізація нелінійного функціоналу нев'язки (2.9) або (2.10) пов'язана з труднощами, спричиненими існуванням декількох локальних мінімумів функції F у просторі значень параметрів. Для подолання цієї проблеми пропонується здійснювати оптимізацію з використанням декількох початкових оцінок значень параметрів. У якості початкових наближень для параметрів тренду вибираються коефіцієнти лінійного тренду, визначені за допомогою МНК.

Вибирається два початкових значення для фази: 0 та π радіан. Це дає можливість підкоректувати фазу апроксимуючої функції як в сторону збільшення, так і в сторону зменшення. Пропонується декілька початкових значень для частоти, вибраних із рівномірною сіткою в інтервалі між мінімальною та максимальною частотою. Емпірично вибрано кількість початкових значень частоти для низькочастотних коливань

(перша ітерація, апроксимація безпосередньо вхідного ряду) $nf = 3$, для середньочастотних (друга та подальші ітерації) $nf = 5$.

Висновки до розділу 2

Запропоновано новий алгоритм прогнозування часових рядів на основі виявлення частотного спектру. Експериментальна апробація показала ефективність запропонованого методу для прогнозування рядів фінансово-економічної динаміки. У порівнянні з відомими класичними підходами, дана технологія дає можливість передбачати моменти часу, в яких можлива зміна тенденції часового ряду, що є важливим у роботі трейдерів на фінансових ринках. При високій інформативності отриманих прогнозів, про які говорять відповідання локальних мінімумів часових рядів прогнозу та реального продовження, класичні критерії оцінювання точності прогнозу (MSE, MAPE) не дають можливості виокремити дану перевагу методу у порівнянні з іншими методами прогнозування. Тому актуальною є розробка нових критеріїв точності прогнозів, які, по-перше, адекватно відобразили б співпадання точок зміни тренду при прогнозуванні, а по-друге, дозволили б спростити задачу пошуку оптимальних параметрів моделі. Похибки прогнозів даного методу можуть бути пояснені наступними моментами: 1) наявністю фінансової кризи на проміжку навчання або прогнозування; 2) похибками моделі частот, що спричиняють невідповідність екстремумів; 3) похибками амплітуд коливань, які не є сталими і можуть збільшуватись під час кризи. Для подолання проблем, які були помічені, пропонується здійснення пошуку оптимальних коефіцієнтів одночасно за двома гармоніками (модель (2.12) при $m=2$) з метою отримати більш складні їх комбінації (накладання частот, “биття”), що має покращити прогнозні можливості запропонованого методу.

РОЗДІЛ 3

ПОБУДОВА ІНФОРМАЦІЙНОЇ СИСТЕМИ ПРОГНОЗУВАННЯ

В даному розділі розглядаються особливості програмної реалізації алгоритмів прогнозування, технології створення програм та архітектура інформаційної системи для прогнозування фондових індексів.

3.1. Аналіз середовищ розробки програмного забезпечення для прогнозування

В даному підрозділі проводиться огляд та аналіз середовищ розробки програмного забезпечення та систем комп'ютерної математики для реалізації алгоритмів, запропонованих в розділі 2. Спочатку розглянемо вимоги до реалізації двох запропонованих алгоритмів прогнозування, оснований на складних ланцюгах Маркова (СЛМ) та на дискретному Фур'є-продовженні (ДФП). Даний підрозділ присвячено особливостям систем розробки програмного забезпечення для реалізації вищерозглянутих методів та алгоритмів на мовах програмування.

Доступні середовища розробки можна умовно поділити на системи розробки програмного забезпечення (ПЗ) загального призначення (Microsoft Visual Studio, GCC, Java, Delphi, Python) та системи математичних розрахунків, які дозволяють розробляти програмні додатки (Matlab, Sage, Scilab, Octave, R).

Розглянемо більш детально дані системи розробки програмного забезпечення.

Таблиця 3.1

Вимоги до середовища розробки

Вимога	СЛМ	ДФП
Імпорт/експорт часового ряду (робота з структурованими файлами, базами даних, веб-сторінками з даними)	+	+
Виконання однотипних операцій з елементами масиву або сусідніми елементами	+	+
Швидка вставка елемента чи послідовності елементів в масив	+	-
Швидке вилучення послідовності елементів з масиву	+	-
Швидке сортування елементів масиву	+	-
Швидкий пошук елементів чи послідовностей елементів масиву	+	-
Оптимізація функції, заданої алгоритмом	-	+

Таблиця 3.3

Середовища програмування

Система розробки ПЗ	ліцензія	мова програмування	математичні бібліотеки	паралельне програмування
Visual Studio 2010	-	C++, C# , Basic, Pascal	Зовнішні (Lapack, Blas та ін)	зовнішні бібліотеки MPI, OpenMP
GNU GCC	GNU	C,C++, Fortran, free pascal	Зовнішні (Lapack, Blas та ін)	Зовнішні бібліотека MPI, OpenMP
Python	GNU	Python	+ (Scipy, Scientific Python)	+
Sage	GNU	Python, Matlab/octave, Maxima, R	+	+
Matlab	-	matlab, Matlab/octave	+	+
Octave	GNU	Matlab/octave	+	+

Системи загального призначення переважно є компіляторами, деякі з них є інтерпретаторами (Python). Для математичних розрахунків, враховуючи значний об'єм спеціалізованих бібліотек та роботу в режимі інтерпретатора, перевага надається системам комп'ютерної математики (Matlab, Octave, Sage), що дозволяє в більш короткі строки реалізувати методи у вигляді програм. У випадку необхідності більшої швидкодії, необхідно реалізовувати алгоритм за допомогою середньорівневих компіляторів з використанням функцій математичних бібліотек (LaPack, Boost та інші). Враховуючи досвід програмування виконавця, та наявністю вільних ліценцій, вибрану систему Octave, яка є інтерпретатором мови, аналогічної до мови в системі matlab. Тому розроблена система може бути використана як у комерційній системі matlab, так і у системі Octave, яка розповсюджується з вільною ліцензією GNU.

3.2. Паралельні алгоритми в методі ланцюгів Маркова та програмні засоби їх реалізації

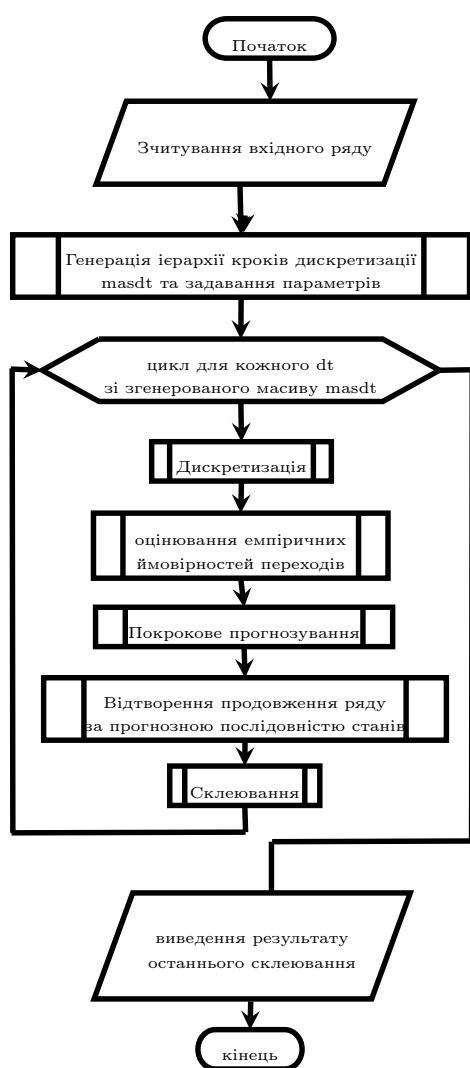
Розроблений метод прогнозування включає наступні процедури: Блок-схема головного алгоритму зображена на рисунку 3.1.

На рис.3.2-3.3 зображено 3 алгоритми класифікації станів.

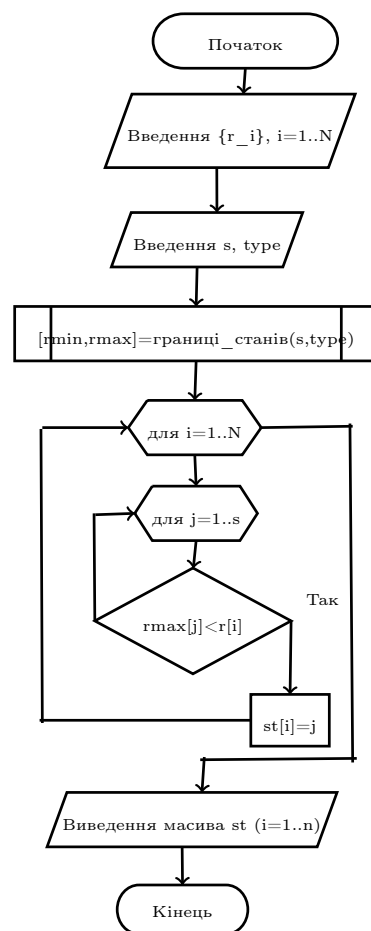
Алгоритм прогнозування на основі ланцюгів Маркова містить такі процедури, які можуть бути оптимізовані з використанням сучасних апаратних та програмних обчислювальних систем:

- процедура перетворення ряду у прирости (формування масиву сусідніх клітинок та обчислення різниці двох рядів).
- процедура класифікації приростів на стани;
- процедура пошуку послідовності станів, еквівалентної даним.

Розглянемо детально алгоритми виконання кожної з процедур та



а)



б)

Рис. 3.1. Основний алгоритм прогнозування (а); Основний алгоритм класифікації прибутковостей (б).

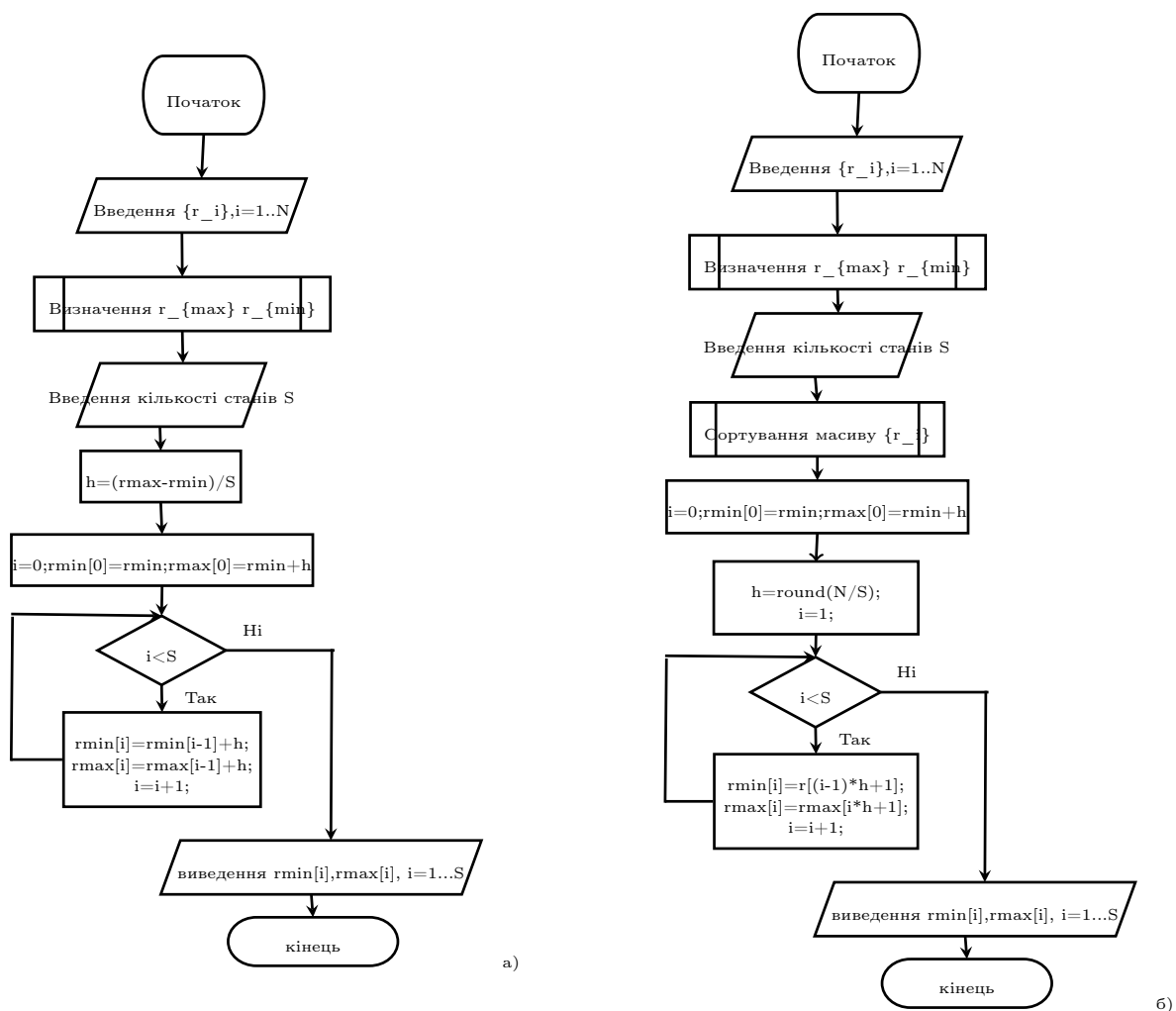


Рис. 3.2. Алгоритм класифікації прибутковостей на основі рівномірності інтервалів (а) та на основі рівномірності представників кожного класу (б)

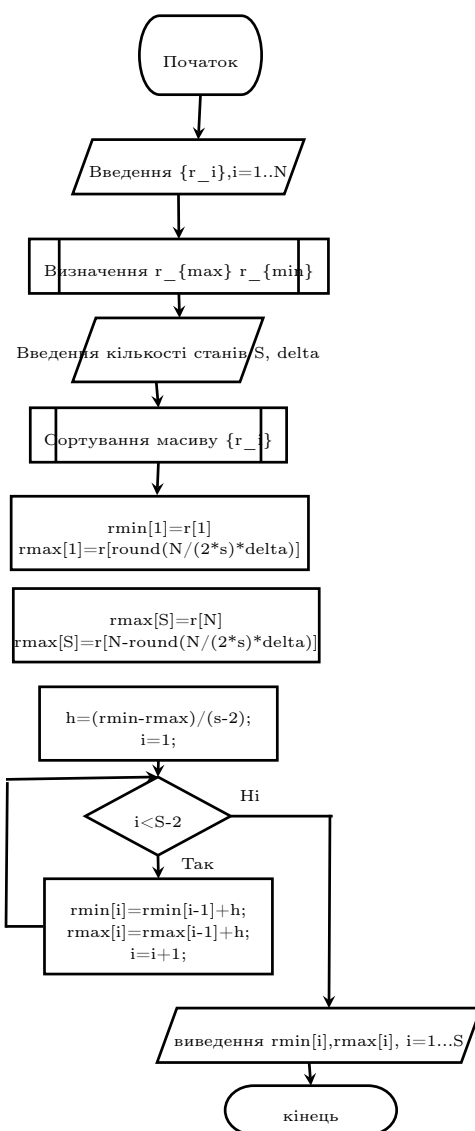


Рис. 3.3. Алгоритм класифікації прибутковостей на основі рівномірної частини “хвоста” розподілу

проаналізуємо можливості їх оптимізації. Програмісту доступні наступні процедури:

- обчислення матричних операцій (додавання, віднімання множення матриць), реалізованих на паралельній архітектурі (векторно-конвейерний принцип, бібліотека MPI);
- поелементні арифметичні дії над елементами вектора (+, -, *, /, порівняння >, <, ==, !=, побітові тощо), які реалізовано як розширеним набором команд процесора (операції над векторами в пам'яті) так і у паралельних архітектурах;
- паралельні та оптимізовані алгоритми сортування масивів даних (наприклад, алгоритм quick sort);
- паралельні алгоритми пошуку (дерева пошуку, зв'язані з масивами даних, хешування умов та швидке повернення результатів);
- паралельні операції з векторами даних, реалізовані через пристрої обробки графічної інформації (GPU).

Оскільки наявність сучасних технологій обчислень у комп'ютерних системах різних користувачів можуть відрізнятись, тому для даної системи прийнято рішення виокремити систему операцій, які можуть бути розпаралелені, розробити загальний інтерфейс їх виконання (усталений формат функцій бібліотеки), та оформити паралельні алгоритми з використанням вибраного базису операцій. В базис операцій включені наступні операції:

- формування масивів (векторів, матриць) однакових елементів (прогресії, специфічні матриці, зокрема реалізовані в Matlab);
- поелементні арифметичні (додавання, віднімання, множення, ділення) та логічні (порівняння, і, або, ні) операції з масивами однакової розмірності;
- поєднання масивів в один (приєднання матриць зліва чи знизу, об'єднання векторів);

- вирізання фрагменту масиву в окремий масив.

Фактично усі ці операції реалізовано в інтерпретованій мові програмування СКМ Matlab ^{??}, з можливістю паралельного їх виконання за допомогою технологій MatlabComputationCluster, MPI та інших. Оскільки система Matlab не є вільно поширюваною, то ми розробили власні реалізації даних процедур на основі різних технологій паралельних обчислень. У випадку послідовної реалізації алгоритму (послідовного виконання виокремлених процедур) дані процедури будуть мати послідовні реалізації, і їх використання, замість прогнозованого зменшення кількості команд, може призвести до його збільшення. Тому послідовний алгоритм, розглянутий у розділі 3, залишається актуальним. Розглянемо паралельну реалізацію процедур алгоритму СЛМ. Процедура обчислення абсолютних або відносних прибутковостей може бути виконана наступним чином (використано мову Matlab):

```

дано: y - вихідний ряд, dt - часовий приріст
n=length(y);
y2=y(dt:n); % формування вектору наступних значень
y1=y(1:n-dt) % формування вектору попередніх значень
r_abs=y2-y1; % обчислення абсолютних прибутковостей (різниці двох
              % сусідніх через dt значень)
r_rel=r_abs./y1;% відносні прибутковості обчислюємо поелементним
               % діленням вектору перших різниць (абсолютних
               % прибутковостей) на вектор попередніх значень.

```

Порівнюючи з послідовною реалізацією (у вигляді циклу), можна стверджувати, що при паралельному виконанні поелементних операцій очікується значне прискорення в часі (що яскраво виявлено при проектуванні алгоритму в системі matlab/octave), у випадку послідовної реалізації кожної операції (віднімання векторів сусідніх значень, ділення) даний алгоритм виконує операцію одного циклу за декілька проходів

циклу (кожна операція - виконання циклу по усім елементам). Тому для реалізації операції - та . / (поелементне ділення) необхідно використовувати швидкі інструкції процесора або паралельність. При їх недоступності (однопроцесорна система) перевагу слід надати послідовній реалізації в один цикл.

Процедура класифікації прибутковостей на дискретні стани. Послідовний варіант даної процедури може бути записаний таким чином (на мові Matlab/octave):

```
% r - прибутковості, r_lim - масив граничних значень прибутковостей
% (межі класів), упорядкованих по зростанню
n=length(r);
for i=1:n % цикл по усім елементам
    for j=1:r_lim % шукаємо перше граничне значення, більше даного.
        if (r(i)>r_lim(j))
            state=j; % номер даного граничного стану - це номер класу,
                    % в який треба віднести
        break;
    end;
end;
state=find(r(i)>r_lim,'first'); % безцикловий варіант на основі
                               %вбудованої в матлаб функції пошуку find
states=[states,state]; % додаємо визначений стан в масив станів.
end;
```

змінна `states` містить необхідну послідовність станів.

Паралельний варіант даного алгоритму може бути реалізовано наступним чином:

```
n=length(r);
states=ones(1,n); % номери станів - масив нулів.
```

% Далі буде використовуватись як лічильник

```

for j=1:length(r_lim)
    statesj=r>r_lim(i);% Поелементна умовна операція. Якщо
        % елемент більший заданої величини (порогу), то
    % результат для даного елементу буде =1,інакше
    % (якщо умова хибна), результат (компонент
    % результуючого вектору) нульовий.
    states=states+statesj; % збільшуються на 1 тільки ті
        % компоненти, які на попередньому кроці дали
    % істину операції порівняння (більші порогу)
end;

```

Таким чином, прирощення відповідних компонентів масиву `states` відбувалось стільки разів, скільки разів прибутковість була вищою порогу j -го стану, тобто дорівнює номеру стану. Дана процедура містить s векторних порівнянь та s векторних додавань, на відміну від максимального $N \cdot s$ порівнянь при пошуку та N додавань до масиву. Але у випадку послідовної реалізації операції “>” та “+”, кількість тактів процесору збільшується, але не перевищує $(N \cdot (s+1))$

Розглянемо операцію пошуку кортежів станів, відповідних даному. Дано: послідовність дискретних станів `states` розмірністю n . Шуканий кортеж `required_states` довжиною r . Необхідно: визначити місця в масиві `states`, які повністю відповідають послідовності `required_states`.

Очевидно, що найпростішим алгоритмом став би цикл пошуку повного співпадання кортежу на вибірці навчання. Але оскільки дана процедура буде викликатись декілька раз для однакової навчальної вибірки `states`, але для різних шуканих станів, то процес пошуку можна оптимізувати. Результати пошуку одного і того ж кортежу при однакових запитах є еквівалентними, тому немає необхідності повторювати пошук повторно. Доречно розробити метод хешування запитів на пошук, зберігати

результати запитів, які найбільш часто викликаються, та при повторних запитах повертати збережений результат.

Для цього використовується хеш-таблиця. Розробимо методику кодування запитів, прономерувавши усі можливі запити від 0 до s^k . Номер запиту можна визначити за формулою:

$$N(\vec{s}) = \sum_{i=1}^S s_i * S^i,$$

де $\vec{s} = (s_1, s_2, \dots, s_r)$ - вектор станів, що шукається, S - загальна кількість дискретних станів. Використовуючи розріджені таблиці, можна зберігати результати запитів (вектор індексів, на яких зустрічається шуканий кортеж). Також можна додатково вести лічильник доступів до елементів розрідженої таблиці. У випадку досягнення розміру таблиці критичного значення, можна видалити записи, які використовувались рідко. Запропонований підхід дозволяє суттєво прискорити прогнозування часових рядів з повторюваними фрагментами.

3.3. Проектування і розробка веб-системи розподіленого прогнозування

В даному параграфі розглядається аналіз вимог та результати проектування веб-орієнтованої системи паралельного та розподіленого прогнозування часових рядів. Така система повинна стати проміжною ланкою між користувачем, який користується веб-сайтом, завантажує ряди або імпортує дані з певних джерел даних для прогнозування та переглядає отримані результати у вигляді графіків, на основі яких приймає рішення. Самі розрахунки ведуться на машинах користувачів, які завантажили систему Octave та скрипт клієнтської програми. Розрахункові клієнти завантажують з серверу дані та вихідні коди методів прогнозування, здійснюють розрахунки прогнозів та відправлення результатів на сервер.

Згідно з підходом об'єктно-орієнтованого проектування [104], процес розробки програмного забезпечення повинен включати наступні кроки:

- аналіз – визначення призначення системи, що розроблюється;
- проектування – визначення підсистем та їх взаємодію;
- реалізація – розробка окремо кожної підсистеми, об'єднання систем в єдине ціле,
- тестування – перевірка роботи системи,
- встановлення – введення системи в дію,
- експлуатація – використання системи.

Розглянемо детально дані процеси розробки бібліотеки програм для прогнозування. На етапі проектування розроблено:

- діаграму прецедентів
- діаграму діяльності
- діаграму класів

Найбільш відомою реалізацією розподілених обчислень, є система BOINC (Berkley Open Infrastructure for Network Computing), яка розроблювалась як платформа проекту по аналізу радіосигналів з космосу з ціллю пошуку позаземних цивілізацій (SETI@HOME). Дана система загальнодоступна на сайті проекту <http://boink.berkley.edu> у вигляді вихідних кодів системи (веб-серверу на php та mysql) та образу передвстановленого серверу на базі ОС Debian Linux з встановленим сервером, компіляторами та шаблонами вихідних кодів для розрахункових програм, а також досить широкого ряду клієнтів для різних ОС (Linux, Windows ME/NT, XP, 7). Перевагою даної системи є розроблений інтерфейс для завантаження даних та програм розрахунку, а також повернення (аплоаду) розрахованих даних на сервер [?].

Оскільки система BOINC розрахована на паралельне виконання компільованих програм, а не інтерпретованих, нами прийняте рішення розробки власної системи паралельного розрахунку прогнозів. Розроблена

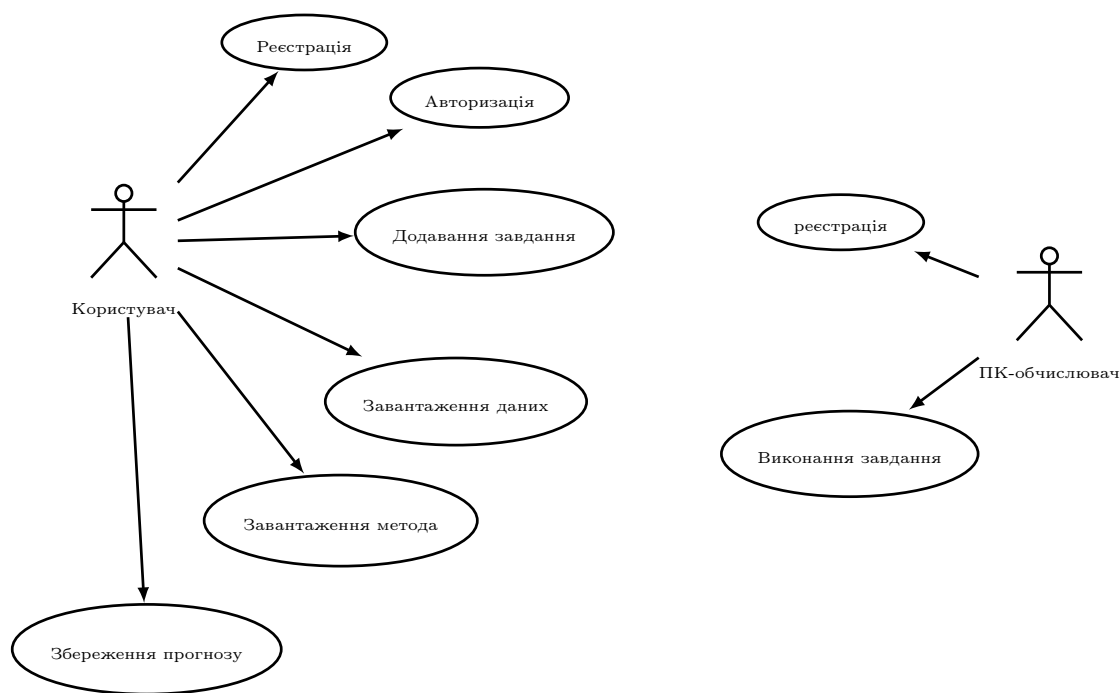


Рис. 3.4. Діаграма прецедентів ІС

нами система орієнтована на виконання програм для інтерпретатора matlab/octave з аналогічним інтерфейсом завантаження/відвантаження даних та програм.

На відміну від системи BOINC, розроблена нами система орієнтована не тільки на розрахунки прогнозів, потрібних адміністратору сервісу. Це є сервіс взаємодопомоги при значному об'ємі розрахунків. Кожен бажаючий може вільно завантажити в систему дані для прогнозування та застосувати до них певний алгоритм (включаючи власне розроблені алгоритми) та отримати обчислювальні потужності інших учасників мережі, які погодились на використання власного машинного часу для розрахунку прогнозів інших учасників.

Діаграма прецедентів системи зображена на рисунку 3.4.

Акторами для даної системи є зареєстровані учасники мережі та адміністратори. Прецеденти, зв'язані з актором “Користувач” наступні:

1. зареєструватись в системі;
2. авторизуватись в системі;

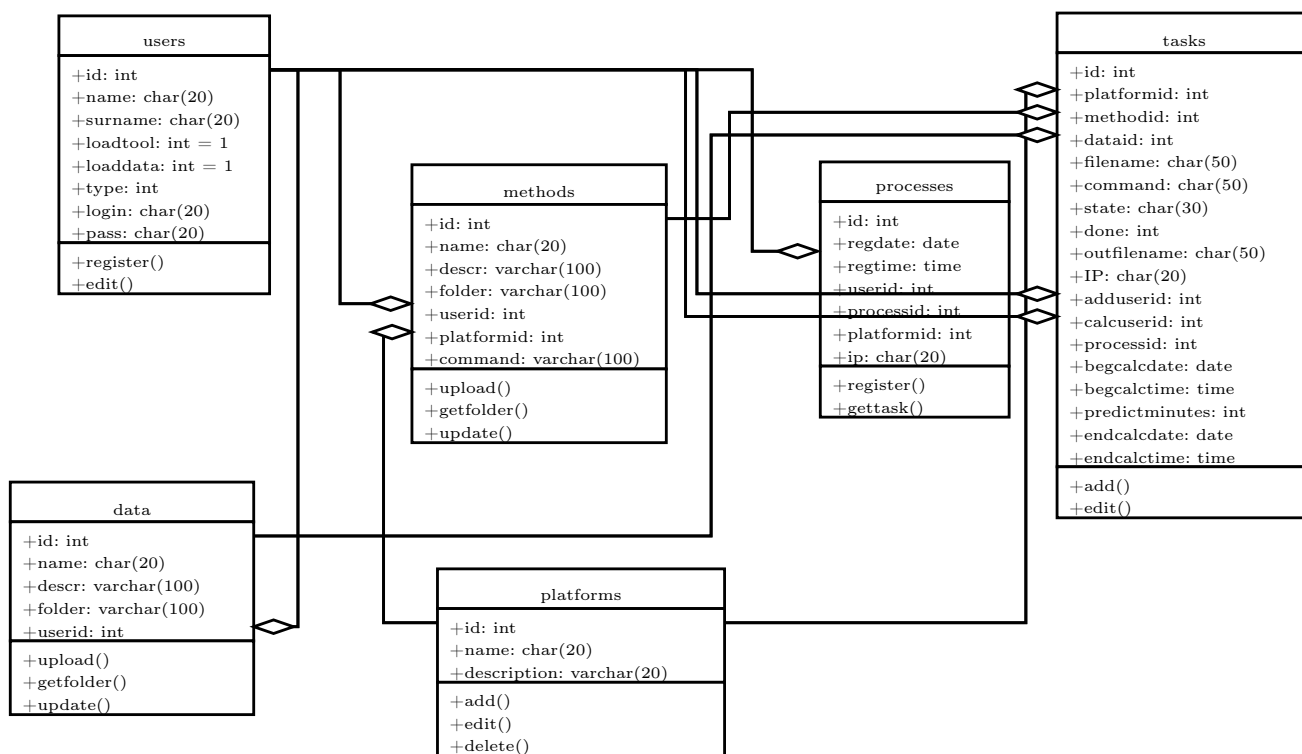


Рис. 3.5. Діаграма класів ІС

3. завантажити/імпортувати прогностні дані ;
4. завантажити алгоритм прогнозування;
5. додати прогнози до групового розрахунку;
6. переглянути/завантажити результати прогнозування.

Прецеденти, зв'язані з актором “ПК-обчислювач” є наступні:

1. зареєструватись в системі (отримати ід, синхронізувати завантажені тулбокси);
2. Завантажити, розрахувати та повернути прогностне завдання;

Діаграма класів системи представлена на рис.3.5. Система має класи: “користувачі” (users), “дані” (data), “методи” (methods), “завдання” (tasks), “розрахункові процеси” (processes). Кожен клас системи відповідає таблиці в базі даних системи.

Розглянемо структуру та призначення кожного класу. Клас

“Користувачі” зберігає інформацію про користувачів, які зареєструвались в системі та має наступні поля:

- id
- Логін
- пароль
- email
- Прізвище, ім'я, по-батькові
- Готовність завантажувати будь-які нові дані
- Готовність завантажити будь-які нові методи.

Клас дані(архіви) призначений для зберігання інформації про завантажений пакет даних для прогнозування та має наступні поля:

- id
- userid (id користувача, який додав пакет)
- name (ім'я пакету даних)
- descr (короткий опис)
- folder (назва папки та архіву)

Клас методи зберігає інформацію про архіви вихідних кодів методів прогнозування:

- id
- name (назва)
- descr (короткий опис)
- folder (папка та ім'я файла)
- userid (id користувача, який додав метод)

Усі дані зберігаються на файловій системі сервера у папках користувачів. У кожного користувача є папки “дані” та “методи”, де у підпапках зберігається файли даних або вихідні коди методів.

При розрахунку, користувач може завантажити або весь пакет, або його частину (новий пакет) або одиничний ряд. Завантажуються дані тільки пакетами. Для оптимізації процесу розрахунку, необхідно зберігати, які

дані завантажив кожен з користувачів та якими методами він погодився розраховувати.

Аналогічно таблиця Методи/користувачі уже вище зазначена. Хоча для розрахунку часто необхідно завантажувати нові методи. Тому необхідно користувачу передбачити ідентифікатор того, що він погоджується завантажити будь-який новий метод або будь-які нові дані. Такі користувачі мають перевагу в правах.

В клас “Користувачі” додані поля “готовність завантажувати нові дані”, “готовність завантажувати нові методи”, значення яких користувач заповнює при реєстрації і визначають, чи відправляти будь-які дані для розрахунку розрахунковим ПК, що належать даному користувачу.

Розглянемо клас “завдання”. Його поля наступні:

- ід
- ід папки даних
- імя файла
- ід метода
- команда (хоча можна сформувати, на основі методу, але зручніше зчитати і викликати)
- імя вихідного файла
- стан (очікується, виконується, помилка)
- відсоток виконання.
- ід користувача, що додав
- ід користувача, що розраховує
- ід клієнтської програми.

Оскільки один користувач може запустити декілька клієнтських програм з однієї ІР-адреси, необхідно мати механізм ідентифікації і розрізнення потоків обчислень між собою. Тому для цього існує таблиця “зареєстровані обчислювальні процеси”. Структура таблиці наступна:

- ід

- час реєстрації
- ід користувача
- ід програми (видається при реєстрації)

Процес запуску клієнта починається з процедури реєстрації. При реєстрації процес отримує ідентифікатор в системі та для нього заповнюється таблиця доступних процесу тулбоксів та даних. При запуску системи дані таблиці порожні, а при виході клієнта з системи відповідні записи очищуються. Можливо, варто таким чином синхронізувати тулбокси (та дані) в системі зі збереженими на клієнтській машині. Такі можливості додано до клієнтів, авторизованих як користувачі з інтерпретатору Matlab, на відміну від розрахункових процесів, не тільки обчислюють завдання з серверу, а можуть додавати власні завдання для розрахунків.

При запиті розрахунковим ПК нових завдань для розрахунку, система здійснює пошук завдання, дані для якого уже завантажені клієнтською програмою та пакет програм (тулбокс) для методу вже встановлено. Якщо жодне завдання, яке очікує розрахунку, не завантажене клієнтом, вибираються ті, які вимагають завантаження тільки тулбоксу, так як тулбокс вимагає менше часу для завантаження та встановлення. Якщо таких завдань у списку очікуваних не знайдено та клієнт погодився на завантаження даних, вибирається завдання, дані для якого вимагаються у більшості завдань, що залишилися.

Дана система при певних умов може працювати нераціонально, зокрема у випадку, коли кількість клієнтів, які завантажили та обчислюють один з пакетів, більша ніж необхідна, іншому пакету може не вистачати розрахункових процесів. Тому актуальною задачею є розробка методів оптимізації розподілу пакетів завдань між користувачами, щоб мінімізувати завантаження нових завдань та максимізувати об'єм виконаної роботи.

Пакети задач повинні бути пропорційно розподілені між усіма учасниками. Необхідно введення певного коефіцієнту часу, або прогнозний час виконання для кожного завдання. Якщо в системі виконуються завдання 2-х пакетів, але сумарно перший пакет може бути обчислений швидше (задачі менші і швидше обчислюються), то кількість копій на 1-шу задачу повинна бути пропорційною обчислювальній долі першого пакету в сумі обчислювального часу усіх доданих завдань. Ідея додати до кожної задачі прогнозний час виконання і на основі задач, що залишились і розрахункових клієнтів, що завантажили даний пакет, прийняти рішення щодо завантаження цієї задачі користувачам з іншої групи (збільшення користувачів, які обчислюють даний пакет), чи виключення певного користувача з групи, які обчислюють даний пакет.

Клієнти архівують файли, які зберігають отримані розрахунки та відправляють їх за допомогою запитів POST протоколу HTTP. Під час виконання завдання можна здійснювати відправку на сервер відсотку виконання завдання для моніторингу виконання. Для передачі та відправлення даних, використовується протокол http. Для запобігання проблемам з розривом зв'язку, файл відправляється пакетами даних з дозаписом отриманого пакету в кінець файлу.

Файл результатів розрахунку зберігаються в підпапці "дані", у кожного користувача при відправці створюється унікальне ім'я файлу, який відправляється на сервер. Після відправки файлу, його вміст розпаковується у папку відповідного завдання. При співпаданні імені (примусовому перерахунку при помилці чи затримці розрахунку), до імені файлів додається номер потоку.

Висновки до розділу 3

Розроблено веб-орієнтовану інформаційну систему для розподіленого обчислення прогнозів. Система складається з web-сервера та множини клієнтів-обчислювачів, які пропонують машинний час своїх персональних ПК для виконання розрахунків, а також клієнтів-користувачів, які через web-інтерфейс мають можливість додати власні дані або власний метод розрахунку на сервер у якості нового завдання для розрахунку. Користувач формує завдання для розрахунку у вигляді zip-архіву вхідних файлів та списку незалежних викликів функцій розрахунку, які будуть застосовані до завантажених даних. На машинах розрахункових клієнтів встановлене програмного забезпечення в системі octave, яке автоматично завантажує дані, оновлену реалізацію методу прогнозування, виконує розрахунки та надсилає їх результати у вигляді zip-архіву.

Реалізація розподіленої системи містить рішення наступних важливих проблем, серед яких: оптимальне розподілення завдань між робочими машинами з огляду на швидкість розрахунку та мінімальність передачі даних та результатів, адаптивна підтримка оптимальності розподілу завдань та перерозподілення у випадках відмови у роботі певних розрахункових машин або включення у роботу нових. Реалізоване повторне відправлення на обчислення задач, які не були обчислені протягом заданого часу. Така реалізація дозволяє обчислювати та публікувати в мережі Internet прогнози фінансових показників в режимі реального часу.

РОЗДІЛ 4

ПРАКТИЧНЕ ВИКОРИСТАННЯ ТА ТЕСТУВАННЯ МЕТОДИКИ СКЛАДНИХ ЛАНЦЮГІВ МАРКОВА

В даному розділі проводиться дослідження таких параметрів методу прогнозування на основі складних ланцюгів Маркова, як кількість станів s та порядок ланцюга Маркова r . Експерименти проводяться на ряді кривих, таких як синусоїда, ряди динамічного хаосу (модель “Хижак-жертва”), а також реальні часові ряди цін акцій та індекси фондових ринків.

4.1. Обґрунтування та експериментальна апробація значень параметрів методу прогнозування, ґрунтованого на складних ланцюгах Маркова

Метод складних ланцюгів Маркова, детально розглянутий у попередньому розділі, має параметри, значення яких впливають на ефективність отриманих прогнозів. Межі значень кожного з параметрів необхідно обґрунтувати теоретичними міркуваннями, а конкретні значення для рядів фінансово-економічної динаміки треба оцінити експериментально. Даний підрозділ присвячено питанню вибору значень параметрів алгоритму, впливу вибору значень на роботу наступних кроків алгоритму та обговоренню результатів експериментальних методів вибору оптимальних значень параметрів та перевірку точності прогнозування розроблених методів.

Розглянемо таблицю, в якій наведено кроки алгоритму СЛМ, перелік параметрів, які задаються на кожному кроці, діапазон їх змін та міри якості виконання кроків.

Таблиця 4.1

Параметри методу СЛМ

Крок алгоритму	Параметри	Міри якості
Вибір множини Δt_k	вибір типу ієрархії (проста чи складна), діапазон Δt_k	кількість значень Δt_k
Обчислення приростів ряду	Абсолютні чи відносні	(у наступних кроках)
Класифікація на стани	1) метод класифікації, 2) Кількість станів s ($2 \leq s \leq 10$), 3) параметри, обчислені з ряду (розмах, дисперсія)	рівномірність представників класів; похибка квантування
Матриця переходів, процес прогнозування	порядок ЛМ r ($1 \leq r \leq 50$), кількість повторів k ($1 \leq k \leq 5$), δ - похибка вибору максимально ймовірного стану ($0 \leq \delta \leq 1$)	ентропія розподілу ймовірностей наступного стану, відсоток відповідності максимально ймовірного стану та реального наступного
Зворотне відновлення прогнозного ряду по станам	Параметри кроку відсутні. Визначаються попередніми параметрами	MSE, MAPE, кореляція прогнозу та реального продовження, співпадання напрямів, відповідність мінімумів

Продовження таблиці 4.1

1	2	3
Склеювання	Визначення значень Δt та параметрами при $\forall \Delta t$	MSE, MAPE та інші, співпадання напрямків, співпадання з прогнозами на інших Δt
Детрендування	Впливає на усі кроки, має параметри, як у апроксимац. функції	усі, зв'язані з якістю прогнозів

Розглянемо детально порядок параметрів, значення яких необхідно вибрати. На першому кроці вибирається множина приростів часу Δt , згідно простої або складної ієрархії (див. розділ 2.3). Оскільки в наступних кроках процес прогнозування буде проводитись для кожного значення Δt_i з вибраної множини, то існує вибір між двома алгоритмами формування множини значень приростів. Це так звані проста (степені числа 2) та складна (степені простих чисел) ієрархії приростів Δt_i . На даному етапі не має можливості оцінити, яка із альтернатив (проста чи складна ієрархія) кращі для якості остаточного прогнозу. Згідно з таблиці 4.1, вибір ієрархії приростів впливає на усі подальші кроки методу. Оскільки результативність вибору множини приростів можна буде оцінити після виконання усіх кроків, то експериментальна перевірка оптимальності ієрархії приростів часу Δt буде обговорена після усіх кроків алгоритму. У якості мір ефективності будуть використовуватись міри відповідності остаточного прогнозного ряду реальному продовженню для кожного приросту Δt .

Наступний крок полягає у обчисленні абсолютних або відносних приростів ряду з вибраним приростом Δt . Зауважимо, що усі інші кроки алгоритму до процедури склеювання, виконуються для кожного Δt з ієрархії приростів. На даному кроці у якості альтернативи для налаштування може вважатись вибір між абсолютними чи відносними приростами. Оскільки якість усіх наступних кроків визначаються вибором абсолютного чи відносного варіанту, то експерименти проведемо з кожною із альтернатив. Міри ефективності кроків для кожної з альтернатив будуть зв'язані з ефективністю остаточного прогнозу. Міри відповідності будуть розглянуті пізніше.

Крок 3 полягає у перетворенні дійсних значень прибутковостей (приростів величин за проміжок часу Δt) у дискретні номери станів. Параметрами даної процедури є кількість станів s ланцюга Маркова,

алгоритм розбиття на стани (рівномірний за частотою, рівномірний за відхиленням, комбінований та комбінований з підсиленням/послабленням хвостів розподілу). Останній метод має параметр для налаштування: степінь підсилення станів на хвостах. Після виконання даного кроку, можна перевірити якість його виконання за допомогою наступних мір:

1. Рівномірність розподілу між станами. Мірами може бути ряд частот кожного стану f та параметри цього розподілу: мінімальна/максимальна частота, розмах частот $R = \max(f) - \min(f)$ чи середня частота стану $\bar{f}$. При нерівномірному розподілі між станами певні стани або зовсім не мають представників, або емпіричні ймовірності переходу у даний стан низькі.
2. похибка квантування, яку можна обчислити, провівши процедуру відновлення ряду без прогнозування по квантованим дискретним станам (див. розділ 2.3). Мірою якості можна обчислити середнє відхилення (MSE) реального та відновленого ряду.

Чисельною мірою похибки квантування візьмемо середнє відхилення вихідного та відновленого ряду (MSE):

У таблиці 4.25 та 4.26 наведено похибки квантування при кодуванню абсолютних та відносних приростах відповідно. В таблицях наведено значення середньої абсолютної похибки при різній кількості станів, довжині відновленого фрагменту, абсолютному та відносному варіанті відхилень. Для більшої достовірності досліджень, взято по 10 фрагментів кожного варіанту довжини, а отримані середні відхилення усереднено.

Як видно з таблиць 4.25, 4.26, 4.3, 4.4, зі збільшенням кількості станів, точність представлення ряду дискретною послідовністю кодів приростів збільшується. Перевагу необхідно надати алгоритму розбиття, при якому відповідна точність досягається меншою кількістю станів s . Тому розумним вибором кількості станів ланцюга Маркова s може бути значення 5-7 станів.

Таблиця 4.3

Квантування та відновлення ряду, приріст 0, MAPE

s	0	1	2	3	4	5
2	51,19682	60,92604	60,85831	44,75592	38,99721	44,36482
3	28,86986	13,45431	27,87102	44,74323	38,99721	28,86351
4	21,2186	22,97351	15,56098	44,74323	23,42777	23,39752
5	25,5024	9,047545	24,55512	44,74323	20,60071	25,92781
6	27,92808	16,17438	20,04214	44,74323	19,6353	29,80099
7	14,10041	8,034824	24,44982	44,74323	22,05435	17,65707
8	14,02937	12,47276	6,592984	44,74323	21,09931	17,83914
9	15,34007	7,896742	24,76989	44,74323	19,35046	18,68655
10	16,07298	9,956271	4,284923	44,74323	18,33006	18,33006
сер. MAPE	23,80651	17,8818	23,2206	44,7446	24,7214	24,9853

У подальших експериментах використовуємо значення s з вищезазначеного діапазону.

Таблиця 4.4

Квантування та відновлення ряду, приріст 1, MAPE

s	0	1	2	3	4	5
2	8,351152	9,822545	9,822297	11,27528	9,878276	8,351152
3	7,444962	8,267926	7,007877	11,27528	9,878276	7,444979
4	5,53369	6,654425	7,612247	11,27526	7,775162	5,67385
5	5,376546	6,380035	7,655186	11,27526	6,422894	6,000041
6	6,175493	7,472644	6,390852	11,27526	6,669136	6,148885
7	6,516722	6,534916	9,634646	11,27523	6,566568	6,114139
8	5,347912	6,691093	5,994673	11,27523	5,866327	5,501067
9	5,425704	5,074064	9,58121	11,27523	6,021554	5,718653
10	6,094001	5,141684	6,080155	11,27523	6,030135	6,030135
сер. MAPE	6,2518	6,8933	7,7532	11,2752	7,2343	6,3314

У попередньому експерименті використовувалась статистика усіх прибутковостей (без розрідження ряду). Близькі за часом прибутковості приймають близькі один до одного значення. Це може спотворити статистику, за якою здійснюється класифікація на стани. Наступний експеримент є аналогічним, але для кожного Δt ряд розріджується, що, можливо, дозволить отримати більш адекватний розподіл. Даний експеримент дозволяє порівняти ці два підходи і зробити висновки про метод, який дає кращі результати відновлення ряду. Результати аналогічного експерименту, але з розріджуванням даних, наведено у таблицях 4.27, 4.28, 4.5 та 4.6.

Таблиця 4.5

Квантування та відновлення ряду, приріст 0, MAPE

s	0	1	2	3	4	5
2	48,31923	59,11249	53,95231	44,17274	26,9297	33,43345
3	28,2323	12,34433	24,0412	44,14046	26,9297	28,66057
4	21,99067	25,83231	23,9021	44,14046	12,18633	23,78298
5	19,42113	8,169812	14,19924	44,14046	8,544357	20,48138
6	13,36357	16,03527	55,64854	44,14046	9,127673	15,86139
7	13,30016	8,225185	19,83945	44,14046	9,396903	15,24754
8	6,894006	12,64329	59,2735	44,14046	7,050711	16,79432
9	6,158324	7,029617	26,18688	44,14046	7,740759	9,433507
10	5,212219	8,679105	59,22746	44,14046	6,516853	6,516853
Сер. MAPE	18,099	17,5635	37,3634	44,144	12,7137	18,9124

З даних таблиць видно, що розрідження дозволяє покращити кодування станів, а відповідно і представлення часового ряду у вигляді послідовності дискретизованих станів. Значення відповідних похибок у таблицях 4.27 - 4.6 є меншими за значення похибок в таблицях 4.25 - 4.4 Кодування

Таблиця 4.6

Квантування та відновлення ряду, приріст 1, MAPE

s	0	1	2	3	4	5
2	9,711837	8,500112	8,574544	8,557691	7,438382	9,882029
3	6,507566	6,781704	6,433049	8,557691	7,438382	6,822469
4	5,018541	4,543221	6,011251	8,557691	4,613703	4,539122
5	3,926606	4,379425	5,034722	8,557691	3,87616	4,376464
6	4,259182	3,409889	3,68036	8,557691	2,477919	3,789604
7	3,714789	2,48852	6,842816	8,557691	2,578385	2,938373
8	3,365443	2,719408	2,857493	8,557691	2,430114	3,069376
9	3,091174	1,570888	6,707816	8,557691	2,010078	2,480251
10	2,388033	2,245349	2,838221	8,557691	2,157872	2,157872
Сер. MAPE	4,6648	4,07	5,4423	8,5577	3,8912	4,4506

у вигляді абсолютних прибутковостей демонструє меншу похибку у більшості експериментів.

Крок 4 полягає у моделюванні послідовності дискретних станів за допомогою складного ланцюга Маркова. Параметром, що підлягає дослідженню, є порядок ланцюгу Маркова r . Оскільки для прийняття достовірного рішення щодо наступного стану, необхідна достатньо велика ймовірність переходу в даний стан. Тому пропонується ще один параметр: k - мінімальна кількість переходів, яка вважається достатньою для прийняття у якості наступного стану.

Зі співвідношення 2.6 та емпірично визначеного інтервалу для параметру $s \in [5..7]$, впливає інтервал для параметра r при відомих s та N :

$$r_{max} = e^{\ln(\frac{N}{k}) \frac{1}{s_{min}+1}}$$

$$r_{min} = e^{\ln(\frac{N}{k}) \frac{1}{s_{max}+1}} \quad (4.1)$$

При середній довжині навчальної вибірки для індексів фондових ринків, $N = 5000$, та кількості представників кожного переходу $k \in [1..1000]$, діапазон зміни значень параметру r буде $1 \leq r \leq 5$. Така коротка пам'ять обумовлена закладеною у співвідношення 2.6 вимогою рівномірно заповнити усі можливі переходи між станами. Але реальні часові ряди є певними реалізаціями випадкового процесу, в яких деякі комбінації зустрічаються рідко, але спричинюють суттєві зміни ряду, що прогнозується. Така властивість пов'язана з невідповідністю ряду прибутковостей нормальному закону розподілу, наявність так званих “важких хвостів” [7]. Тому висунуто гіпотезу, що повторення однієї комбінації всього 2 рази достатньо, щоб вибрати її продовження як прогноз. Тоді, зафіксувавши мінімальну кількість повторів $k = 2$, запропоновано метод автоматичного вибору порядку ланцюгу Маркова в залежності від довжини послідовностей станів СЛМ, які повторювались більш ніж k разів. Для запобігання копіювання певного фрагмента у прогноз, автоматично вибраний параметр r пропонується обмежити зверху. Таким чином, експериментальному апробуванню підлягають параметри r_{max} та k .

Експериментальні результати по прогнозуванні з $k = 2$ показують значення середнього порядку ланцюга Маркова на рівні $r \approx 50$, що відповідає нулю автокореляційної функції для модулів прибутковості. Це дає можливість стверджувати, що оптимальна довжина пам'яті для рядів фондових ринків розвинених країн (США, ЄС, Японія та ін.) є порядку 50 днів.

Для кількісного оцінювання якості виконання кроку 4, крім мір оцінки остаточних прогнозів, можна використовувати ентропію розподілу ймовірностей наступного кроку. За основу взята ентропія Шеннона:

$$E = - \sum_{i=1}^s p_i \cdot \ln p_i \quad (4.2)$$

Дана міра яскраво відображає випадки невизначеності, коли декілька станів мають високу ймовірність. Як уже було зазначено у розділі 2, сусідні стани з максимальною ймовірністю можуть бути мисленно об'єднані в кластери, що дає можливість запропонувати правила подолання таких невизначеностей. Ентропія Шеннона не відрізняє місце масимально ймовірних станів у розподілі, але дає можливість, не проводячи процедуру прогнозування до кінця, оцінити ступінь невизначеності при виборі прогнозного стану.

Ще однією мірою якості вибору конфігурації вищерозглянутих параметрів є узгодженість стану з максимальною ймовірністю та реального наступного стану. Оцінити таку відповідність можна як на навчальній вибірці без використання нової перевірконої вибірки, так і на новій вибірці. Перший випадок може бути доречним на етапі калібровки параметрів для нового ряду. При незадовільних значеннях мір відповідності, користувач мусить прийняти рішення щодо відсутності у вихідному часовому ряді вираженої марківської залежності.

У якості міри узгодження може бути взята рангова кореляція Спірмена:

$$r = 1 - \frac{6 \sum d^2}{n(n^2 - 1)} \quad (4.3)$$

де $\sum d^2$ - сума квадратів різниць рангів, n - число парних порівнянь. У якості рангів взятий порядок стану s_i у ряді станів, упорядкованим за ймовірністю p_i .

Останній крок для вибраного інтервалу дискретизації Δt , полягає у відновленні прогнозного ряду з ряду дискретизованих станів складного ланцюга Маркова. Даний крок, так як і процедура склеювання, не містить нових параметрів, а повністю визначається параметрами попередніх

кроків. Фактично процедура відновлення та склеювання видають ряд спрогнозованих значень. На основі прогнозного ряду можна обчислити міри якості отриманого прогнозу за допомогою наступних відомих мір: середній квадрат відхилень (Mean Squared Error, MSE):

$$MSE = \frac{1}{N} \sqrt{\sum_{i=1}^N (y_i - y(x_i))^2}, \quad (4.4)$$

середня абсолютна похибка (Mean Absolute Error, MAE):

$$MAE = \frac{1}{N} \sum_{i=1}^N |y_i - y(x_i)|, \quad (4.5)$$

середня абсолютна похибка у відсотках (Mean Absolute Percent Error, MAPE):

$$MAPE = \frac{1}{N} \sum_{i=1}^N \frac{|y_i - y(x_i)|}{\bar{y}} \cdot 100\% \quad (4.6)$$

та коефіцієнт невідповідності Г. Тейла:

$$U = \sqrt{\frac{\sum_{i=1}^N (y_i - y(x_i))^2}{\sum_{i=1}^N (y_i - y(x_i))}}. \quad (4.7)$$

Дані міри відображують тільки відхилення фактичних значень від спрогнозованих. Але прогнозний ряд є основою прийняття інвестиційних рішень, і трейдеру більш важливо прогнозувати зміни чи переломи трендів, моменти локальних екстремумів, які можуть бути представлені на прогнозній кривій, не зважаючи на відхилення від фактичних значень. Тому актуальною задачею є розробка нових мір якості прогнозів, які дали б змогу відобразити не тільки відхилення прогнозу від фактичного значення у певний момент, але і відхилення прогнозованого часу локального

екстремуму від часу фактичного екстремуму. Для цього пропонуються дві процедури.

Ідея першої полягає у оцінюванні похибок напрямку прогнозу.

Друга ідея базується на визначенні моментів часу локальних екстремумів на фактичному та прогнозованому ряді, та обчислення похибки між моментами відповідних екстремумів у одиницях часу.

4.2. Тестування прогнозу фінансового часового ряду методом складних ланцюгів Маркова

В даному підрозділі наводяться результати експерименту по застосуванню параметрів, одержаних в результаті проведених експериментів у процесі прогнозування реального фінансового ряду. Наводиться порівняння ефективності розробленого методу прогнозування з відомими методами:

1. Авторегресійна модель ковзного середнього
2. Нейронні мережі
3. МГУА

У даному експерименті використано часові ряди індексів фондового ринку США Dow Jones Industrial Average (DJIA), та індексу S&P 500. Фактично дані часові ряди є середнім зваженим котирувань акцій найбільш важливих акцій ринку (30 індустріальних компаній у індексі DJIA та 500 компаній у індексі S&P 500). Оскільки ціни акцій на розвинених фондових ринків є корельованими, дані індекси використовуються учасниками ринку як показники загальної тенденції на ринку. Також на сучасних ринках (включаючи російський ММВБ-РТС та український UX) існують такі інвестиційні інструменти як ф'ючерси на індекс, які дозволяють замість торгівлі пакетами акцій, взятими з пропорціями, відповідними індексному кошику, здійснювати спекулятивну торгівлю. Оскільки після відкриття

ф'ючерного контракту (на купівлю або продаж) обов'язково необхідно здійснювати закриття позиції, тобто виконувати протилежну операцію, то для укладання ф'ючерного контракту не обов'язково мати повну вартість, відповідну значенням індексу, використовується кредитне плече (для ф'ючерса на український індекс UX це $1/4$). З огляду на вищесказане, задача прогнозування індексів фондових ринків є дуже актуальним для інвесторів.

Розглянемо результати прогнозування світових фондових індексів на основі моделей авторегресії та ковзного середнього. Прогноз виконувався за допомогою моделі $ARMA(p, q)$, де параметри p та q вибирались з інтервалу ($5 \leq p \leq 10$ та $0 \leq q \leq 2$). Кожна комбінація значень параметрів проходила попереднє тестування на перевіірочній вибірці. Для побудови остаточного прогнозу вибиралась пара значень p та q , яка показала найменшу похибку на перевіірочній вибірці. Графіки деяких прогнозів моделями $ARIMA(p, 1, q)$ зображено на наступних рисунках.

На рисунках 4.15 у додатку 2 наведено характерні режими, які демонструються авторегресійною моделлю...

В наступній таблиці 4.28 наведено похибки остаточних прогнозів та прогнозів методики СЛМ за цей же період (MSE).

Таблиця 4.7

Похибки прогнозування методом ARMA

Індекс	довжина навч вибірки	початк. точка	кінцева точка	похибка ARMA
1	2	3	4	5
DJI	10500	1	10501	$1.17 \cdot 10^4$
DJI	10500	1795	12295	$9.5 \cdot 10^4$

Продовження таблиці 4.28

1	2	3	4	5
DJI	10500	3589	14089	$5.46 \cdot 10^5$
DJI	10500	5383	15883	$4.92 \cdot 10^6$
DJI	10500	7177	17677	$3.95 \cdot 10^6$
DJI	10500	8971	19471	$3.6 \cdot 10^6$
DJI	15000	1	15001	$2.62 \cdot 10^6$
DJI	15000	895	15895	$5.98 \cdot 10^6$
DJI	15000	1789	16789	$2.39 \cdot 10^7$
DJI	15000	2683	17683	$2.95 \cdot 10^6$
DJI	15000	3577	18577	$2.0 \cdot 10^6$
DJI	2000	1	2001	$9.77 \cdot 10^3$
DJI	2000	3495	5495	$4.78 \cdot 10^3$
DJI	2000	6989	8989	$2.85 \cdot 10^4$
DJI	2000	10483	12483	$4.62 \cdot 10^3$
DJI	2000	13977	15977	$1.39 \cdot 10^5$
DJI	2000	17471	19471	$3.96 \cdot 10^6$
DJI	3000	1	3001	$3.2 \cdot 10^3$
DJI	3000	3295	6295	$2.16 \cdot 10^4$
DJI	3000	6589	9589	$3.54 \cdot 10^4$
DJI	3000	9883	12883	$3.72 \cdot 10^4$
DJI	3000	13177	16177	$1.41 \cdot 10^6$
DJI	3000	13177	16177	$1.41 \cdot 10^6$
DJI	3000	16471	19471	$4.57 \cdot 10^6$
DJI	4500	1	4501	$1.64 \cdot 10^2$
DJI	4500	2995	7495	$8.36 \cdot 10^4$
DJI	4500	5989	10489	$2.28 \cdot 10^4$
DJI	4500	8983	13483	$2.34 \cdot 10^5$

Продовження таблиці 4.28

1	2	3	4	5
DJI	4500	11977	16477	$4.24 \cdot 10^6$
DJI	4500	14971	19471	$7.49 \cdot 10^6$
DJI	7000	1	7001	$2.67 \cdot 10^4$
DJI	7000	2495	9495	$1.14 \cdot 10^5$
DJI	7000	4989	11989	$3.38 \cdot 10^4$
DJI	7000	7483	14483	$9.3 \cdot 10^5$
DJI	7000	9977	16977	$4.0 \cdot 10^7$
DJI	7000	12471	19471	$3.93 \cdot 10^6$

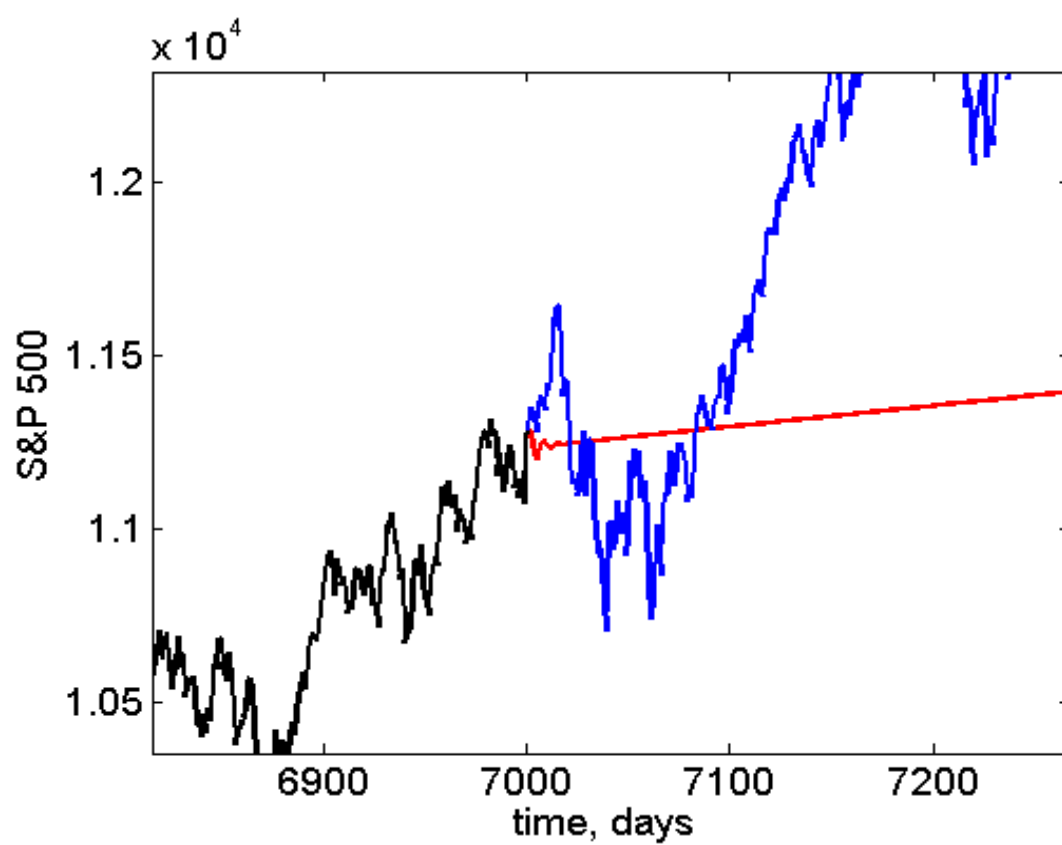


Рис. 4.1. Прогноз індексу Dow Jones Industrial моделлю ARMA(5,1)

Динаміка системи задана послідовністю значень $y_i = y(t_i)$, ($i = 1..N$) показника, що прогнозується (у даному випадку - денні значення індексу фінансового ринка), вимірюного з постійним кроком дискретизації Δt . (у даному експерименті - один день). Необхідно оцінити величини майбутніх значень $y_{N+1}, y_{N+2}, \dots, y_{N+NP}$, базуючись на закономірностях, виявлених вибраною прогновною моделлю.

Для забезпечення статистичної стаціонарності та повторюваності навчальної вибірки, були обчислені прирости (прибутковості) часового ряду:

$$r_t = \frac{y(t) - y(t - \Delta t_{new})}{y(t)}, \quad (4.8)$$

де r_t - значення відносних приростів, в подальшому використовується як вхідний сигнал для нейронної мережі. Δt_{new}

Для прогнозування вибрано лагову залежність виду:

$$r_{t+\Delta t_{new}} = f(r_t, r_{t-\Delta t_{new}}, \dots, y_{t-m\Delta t_{new}}), \quad (4.9)$$

де r_t - величина показника, що прогнозується в момент часу t (фактично часовий ряд відносних приростів); m - довжина “пам'яті” ряду, які використовуються при прогнозуванні; Δt_{new} - новий крок дискретизації (змінний параметр, є кратним кроку дискретизації вихідного ряду, фактично використовується розрідження часового ряду).

Співвідношення (4.9) було використано ітераційно, додаючи спрогнозоване значення в кінець вихідного часового ряду, таким чином реалізовано багатокрокове прогнозування часового ряду. Здійснено 1000 ітерацій, що відповідає прогнозуванню на 4 роки вперед. Отриманий ряд спрогнозованих відносних приростів піддавався процедурі відновлення:

$$y_{t+\Delta t_{new}} = y(t)(1 + r_{t-\Delta t_{new}}), \quad (4.10)$$

Середні абсолютні (MAE) та квадратичні (MSE) відхилення реального продовження та прогнозного часового ряду використовувались як міри якості прогнозу.

Оскільки основна ціль даного експерименту - порівняти прогностичні можливості моделей, що базуються на нейронних мережах з методом складних ланцюгів Маркова, то довжина пам'яті вибрана відповідно до довжини пам'яті, що використовувались у методі складних ланцюгів Маркова ($m = 30$ для Δt_{new} від 1 до 32 днів; $m = 20$ для Δt_{new} 64 дня; $m = 10$ для $\Delta t_{new} = 128$ та $m = 5$ для $\Delta t_{new} = 256$). Дані значення дуже близькі до значення нуля автокореляційної функції для модулів відносних прибутковостей, що є додатковим підтвердженням правильності вибору параметру m .

Оскільки у методиці складних ланцюгів Маркова використано методику “склеювання”, щоб використати прогнози на усіх частотних рівнях (при інтервалах дискретизації $\Delta t_{new} = 1, 2, 4, 8, 16, 32, 64, 128, 256$) а класична лагова залежність (4.9) не має такої процедури, то проведено експерименти з використанням усіх можливих значень Δt_{new} та проведено аналіз результатів. Виявилось, що при $\Delta t_{new} = 32$ як методика складних ланцюгів Маркова, так і нейронних мереж дають найкращі результати. Результати експерименту та його інтерпретація наведена далі у даній роботі.

Для побудови моделі відображення $f(\dots)$ використовується нейронні мережі різної архітектури.

Для проведення експерименту використовувався Neural Networks Toolbox системи Matlab. За даними навчальної вибірки було сформовано множину векторів, яка піддана апробації усіма доступними у Neural Networks Toolbox архітектурами нейронних мереж. Після візуального аналізу отриманих результатів було відкинуто ті архітектури, які не здатні розв'язувати дану задачу. Було відібрано для подальших експериментів наступні архітектури:

1. Cascade Back Propagation
2. Elman Back Propagation
3. Feed Forward Back Propagation
4. Linear Layer (design)
5. NARX (delays)
6. Perceptron
7. Radial Basis Fewer neurons

Саме така нумерація актуальна для таблиць з похибками прогнозування, наведеними нижче.

У всіх мереж вибрано 1 вхідний шар, 1 прихований та 1 вихідний нейрон з лінійною функцією активації. Для попереднього відбору архітектури використано 10 нейронів прихованого шару.

Експерименти для обґрунтування вибору параметрів методу

Для експериментів з вихідного часового ряду індексів DJIA та S&P 500 вибрано по 5 фрагментів у якості навчальної вибірки різної довжини (2000, 3000 та 4500 днів. Всього фрагментів - 15 для DJIA та 15 для S&P 500). Здійснено прогнозування кожного з фрагментів з $\Delta t_{new} = 1, 2, 4, 8, 16, 32, 64, 128$ та 256. Довжина прогнозу - 1000 днів. Результати прогнозування наведено у наступній таблиці:

Таблиця 4.9

Похибки прогнозування (MSE) при різному інтервалі дискретизації Δt_{new}

type	1	2	4	8	16	32	64	128	256	Макс.
type	1	2	4	8	16	32	64	128	256	Макс.
1	6,61E+07	1,93E+245	4,32E+07	1,27E+49	2,23E+07	5,40E+06	1,79E+07	8,33E+06	1,24E+07	1,93E+245
2	5,98E+06	6,77E+06	2,87E+23	2,93E+07	2,62E+07	3,96E+07	9,85E+08	5,43E+07	3,03E+13	2,87E+23
3	1,42E+35	1,27E+15	4,66E+198	2,43E+08	2,32E+114	2,85E+07	2,24E+08	1,04E+07	2,00E+07	4,66E+198
4	1,35E+68	9,92E+150	4,25E+66	7,32E+09	9,58E+15	1,79E+07	4,41E+08	9,76E+06	2,91E+07	9,92E+150
5	1,89E+09	7,47E+15	3,36E+08	3,26E+10	1,07E+07	3,24E+07	2,13E+08	4,37E+06	1,08E+07	7,47E+15
6	1,97E+07	1,91E+07	1,91E+07	1,64E+07	1,44E+07	1,31E+07	5,19E+08	8,75E+06	5,85E+07	5,19E+08
7	2,62E+07	2,63E+07	2,01E+07	1,99E+07	1,51E+07	1,50E+07	5,32E+08	1,10E+07	5,74E+07	5,32E+08
Макс.	1,35E+68	1,93E+245	4,66E+198	1,27E+49	2,32E+114	3,96E+07	9,85E+08	5,43E+07	3,03E+13	1,93E+245

По рядкам даної таблиці - типи архітектур нейромережі (номерація відповідає номерації рисунків у попередньому розділі). По стовбцям

- вибраний новий інтервал дискретизації у днях. Значення таблиці - максимальне відхилення продовження. З даної таблиці видно, що найкращі результати відповідають інтервалу дискретизації $\Delta t_{new} = 32$. Найкращі прогнози демонструють архітектури нейромережі № 6 (Perceptron) та № 7 (Radial Basis Fewer neurons).

У наступних таблицях результати погруповано згідно довжини навчальної вибірки та архітектури нейронної мережі. Проміжок дискретизації вибрано 32, згідно висновків з попередньої таблиці.

Таблиця 4.10

Похибки прогнозування (MSE) з різною довжиною навчальної вибірки:

S&P 500

Довж. навч. вибірки	1	2	3	4	5	6	7	Макс.
2000	60 795	164 044	142 835	124 465	186 193	185 202	194 202	194 202
3000	40 430	196 631	170 710	149 801	136 433	189 153	186 662	196 631
4500	814 093	334 362	490 152	3 981 184	2 865 029	95 800	112 323	3 981 184
Макс.	814 093	334 362	490 152	3 981 184	2 865 029	189 153	194 202	3 981 184

Таблиця 4.11

Похибки прогнозування (MSE) з різною довжиною навчальної вибірки:

DJIA

Довж. навч. вибірки	1	2	3	4	5	6	7	Макс.
2000	3 640 634	6 641 535	6 057 511	3 644 709	5 775 623	6 679 388	6 589 994	6 679 388
3000	4 472 468	3 653 266	19 303 570	4 984 013	15 241 700	7 175 845	7 167 810	19 303 570
4500	5 396 986	39 608 450	28 471 140	17 941 570	32 363 030	13 137 780	15 048 810	39 608 450
Макс.	5 396 986	39 608 450	28 471 140	17 941 570	32 363 030	13 137 780	15 048 810	39 608 450

Наведені вище таблиці містять по стовбцях - архітектури нейромережі, по рядках - довжина навчальної вибірки. У комірках наведено максимальна похибка прогнозу на 1000 кроків. З даного експерименту можна зробити висновки, що найкращі прогнози отримано при довжині навчальної вибірки 2000 з використанням архітектур з номерами 1 (Cascade Back Propagation) та 4 (Linear Layer (design)).

Таблиця 4.12

Похибки прогнозування методу СЛМ та нейронних мереж

Індекс	Довжина навч вибірки	NeuralMSE	MarkovMSE	Абс. значення	Нейр MAPE	Марков MAPE
DJIA	2000	12 232,27	15 407,07	890,85	12,42%	13,93%
DJIA	2000	908 109,60	133 996,19	3 404,10	27,99%	10,75%
DJIA	2000	3 640 634	2 657 171	11 279,00	16,92%	14,45%
DJIA	2000	366,16	836,27	226,36	8,45%	12,78%
DJIA	2000	41 140,54	39 145,89	829,21	24,46%	23,86%
DJIA	3000	700 298,20	346 047,47	3 461,90	24,17%	16,99%
DJIA	3000	4 472 468	8 127 140	11 279,00	18,75%	25,28%
DJIA	3000	25 670,94	9 563,60	279,52	57,32%	34,99%
DJIA	3000	5 315,59	26 032,27	835,13	8,73%	19,32%
DJIA	3000	73 741,43	27 944,20	819,54	33,13%	20,40%
DJIA	4500	3 142 753	2 377 172	3 758,40	47,17%	41,02%
DJIA	4500	5 396 986	5 999 085	11 279,00	20,60%	21,72%
DJIA	4500	15 744,05	8 576,77	522,34	24,02%	17,73%
DJIA	4500	26 969,67	4 388,64	758,49	21,65%	8,73%
DJIA	4500	237 617,30	216 434,67	811,83	60,04%	57,31%
S&P 500	2000	316,21	36,72	39,48	45,04%	15,35%
S&P 500	2000	456,76	141,88	94,25	22,68%	12,64%
S&P 500	2000	194,50	235,38	99,45	14,02%	15,43%
S&P 500	2000	23 598,82	1 137,88	255,33	60,16%	13,21%
S&P 500	2000	60 795,00	133 565,76	665,69	37,04%	54,90%
S&P 500	3000	400,54	73,19	72,64	27,55%	11,78%
S&P 500	3000	901,98	175,18	88,48	33,94%	14,96%
S&P 500	3000	276,39	1 018,53	111,27	14,94%	28,68%
S&P 500	3000	7 592,51	2 532,43	271,15	32,14%	18,56%
S&P 500	3000	40 430,15	25 882,18	843,43	23,84%	19,07%
S&P 500	4500	574,23	121,76	94,17	25,45%	11,72%
S&P 500	4500	414,70	75,22	86,20	23,62%	10,06%

Продовження таблиці 4.12

1	2	3	4	5	6	7
S&P 500	4500	833,21	1 692,19	162,95	17,71%	25,24%
S&P 500	4500	4 695,10	7 177,01	330,20	20,75%	25,66%
S&P 500	4500	814 092,60	60 205,85	1 074,90	83,94%	22,83%

На наступних рисунках зображено прогнозну динаміку та порівняння з реальним продовженням для експериментів з довжиною навчальної вибірки 2000 днів.

З таблиці 4.12 та рисунків 4.16-4.24 видно, що при прогнозуванні на 1000 днів вперед найкращі результати для методу нейронних мереж отримані при довжині вибірки 2000 днів та мережі Cascade Back Propagation із бібліотеки Neural Network Toolbox системи Matlab. Оптимальний інтервал дискретизації для обчислення відносних приростів - 32 дня. У 20 випадках з 30 (67 %) метод складних ланцюгів Маркова демонструють більш точні прогнози, що говорить про перспективність цього нового методу прогнозування часових рядів.

Таблиця 4.14

Похибки прогнозування індексу Dow Jones Industrial методами СЛМ та нейронних мереж

	Markov MSE	1	2	3	4	5	6	7
DJI	4862,81	29825,54	771,8155	419,0937	333,8829	708,146	478,9897	515,4856
DJI	15407,07	5050,792	4603,083	7748,256	6278,986	3579,235	3697,634	3622,063
DJI	133996,19	99226,58	125310	57176,59	150437,5	60513,2	132736,3	159474,1
DJI	2657170,59	3640634	4285888	3289487	3133443	3877453	3818135	4060255
DJI	836,27	343,21	271,7977	177,7498	271,6964	101,2644	248,7703	209,4772
DJI	39145,89	4223,914	5620,258	4370,172	4738,079	3226,095	2362,942	1772,327
DJI	196,85	641,2571	123,276	247,2131	369,4413	392,6245	215,7118	303,4189
DJI	346047,47	139694,1	181015,7	227291	97476,5	148175,4	130119,7	157027,7
DJI	8127140,46	3558192	2895846	6013268	2527864	1595068	7175845	7167810
DJI	9563,6	5461,267	7773,935	3562,303	5484,123	4309,42	11234,14	4566,526
DJI	26032,27	5315,589	3430,66	16557,13	3979,254	19935,65	5558,238	5383,096
DJI	27944,2	14354,19	30557,61	8645,996	12666,35	9136,454	7639,301	7562,249

Продовження таблиці 4.14

1	2	3	4	5	6	7	8	
DJI	18358,54	700,0981	187,3698	155,7703	480,0762	163,0023	225,6405	311,0369
DJI	2377171,9	106450	593890,6	1058077	1191157	1274379	1486320	980308,4
DJI	5999085,19	4897333	4234222	3035355	4724475	2769336	6879028	5460036
DJI	8576,77	9535,564	4304,265	2012,591	3541,091	7738,709	3535,161	3134,458
DJI	4388,64	14522,97	8456,225	10093,76	10846,8	11436,65	8643,565	8008,173
DJI	216434,67	73562,44	131504,7	42296,75	142283,7	19161,62	128344,5	128220,5

Таблиця 4.16

Похибки прогнозування індексу S&P 500 методами СЛМ та нейронних мереж

	Markov MSE	1	2	3	4	5	6	7
GSPC	36,72	76,33231	46,8864	49,58305	42,46866	83,45371	41,87643	31,16411
GSPC	141,88	89,09176	67,92983	127,9769	198,8137	175,6017	240,6194	212,3979
GSPC	235,38	89,35446	452,4533	157,2896	125,3486	109,1961	89,51768	89,55271
GSPC	1137,88	1438,175	1184,36	727,8405	777,1543	852,8376	537,3756	521,8377
GSPC	133565,76	18409,5	46718,79	7035,171	63977,94	51162,44	70696,23	46883,95
GSPC	73,19	88,60546	86,91811	61,01672	37,0717	77,47533	68,48738	63,45674
GSPC	175,18	235,427	109,4156	219,5497	125,0683	72,31041	139,8682	149,3686
GSPC	1018,53	189,6797	214,0317	188,6304	399,4031	100,1906	234,7954	186,769
GSPC	2532,43	752,8921	540,3072	2567,136	799,5384	2645,999	392,5501	439,1967
GSPC	25882,18	11086,59	17820,59	10803,65	51369,16	8611,761	23459,13	20588,33
GSPC	121,756	195,629	87,72242	68,19237	113,8921	91,08167	273,81	275,2606
GSPC	75,22	397,9204	70,7982	122,5436	226,1052	113,7206	58,91802	60,76534
GSPC	1692,19	96,46894	1562,415	478,9659	347,6168	575,9448	793,6815	722,2358
GSPC	7177,0	981,3162	1421,711	608,6635	263,4561	653,4561	201,1695	195,4168
GSPC	60205,85	126389,2	44060,29	41708,47	30733,04	9448,52	26730,99	26564,83

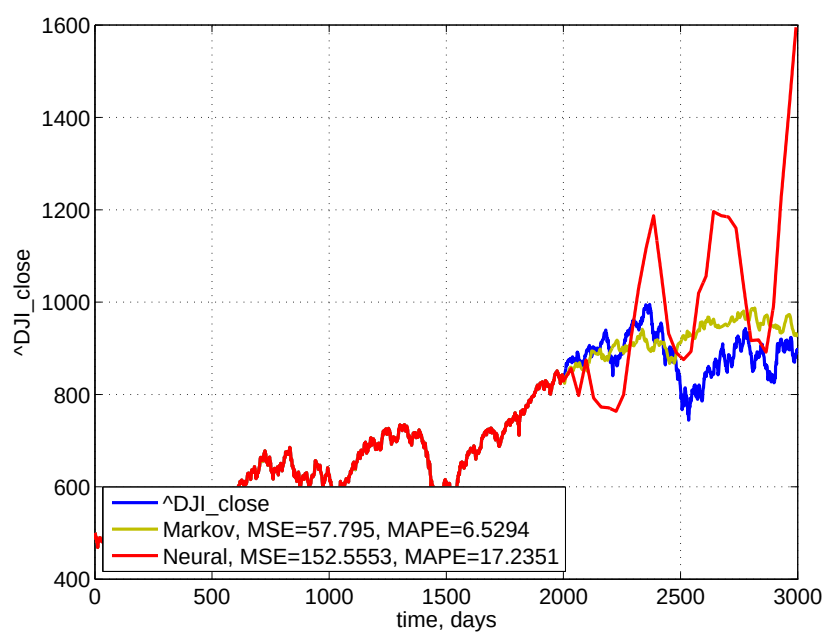


Рис. 4.2. DJIA, learningset: days 6989-8989

Таблиця 4.18

Похибки прогнозування індексу Dow Jones методами СЛМ та нейронних мереж

	Markov MAPE	1	2	3	4	5	6	7
DJ 2000 (1)	29,05%	99,74%	5810,5%	74,79%	$1,4 \cdot 10^5$	6038,19%	51,72%	51,79%
DJ 2000 (2)	13,93%	58,70%	23618,7%	56,21%	32,34%	84,37%	19,34%	18,87%
DJ 2000 (3)	10,75%	29,43%	24641,7%	$3,4937E+03$	106714,00%	247,92%	18,05%	22,02%
DJ 2000 (4)	14,45%	36,39%	8352,86%	84,37%	766655%	78,62%	22,63%	22,37%
DJ 2000 (5)	12,78%	68,33%	100618,3%	7296,89%	$3,3 \cdot 10^5$	34,10%	12,14%	13,92%
DJ 2000 (6)	23,86%	17,95%	221294,6%	24,05%	44,90%	28,60%	14,84%	12,83%
DJ 3000 (1)	5,85%	84,71%	15,24%	391,58%	3016,77%	1151,61%	13,73%	14,17%
DJ 3000 (2)	16,99%	22,86%	24,54%	214,78%	$1,45 \cdot 10^{15}$	21,52%	14,93%	19,40%
DJ 3000 (3)	25,28%	36,81%	48,54%	43,73%	60969,2%	74,56%	34,29%	33,59%

Продовження таблиці 4.18

1	2	3	4	5	6	7	8	
DJ 3000 (4)	34,99%	44,17%	44,61%	51,73%	92,30%	61,94%	43,66%	42,90%
DJ 3000 (5)	19,32%	76,47%	22,41%	53,07%	2570547%	7204,74%	23,79%	23,72%
DJ 3000 (6)	19,96%	41,11%	31,02%	$1,5 \cdot 10^{14}$	39,69%	36,71%	23,99%	25,91%
DJ 4500 (1)	56,45%	48,12%	15,24%	73,67%	821,51%	26,83%	17,72%	19,62%
DJ 4500 (2)	41,02%	51,57%	46,01%	57,28%	28224,74%	1455,21%	40,50%	39,44%
DJ 4500 (3)	21,72%	39,72%	100,72%	66,90%	91,12%	64,13%	77,33%	78,86%
DJ 4500 (4)	17,73%	87,64%	26,49%	28,84%	36,76%	39,80%	16,41%	16,91%
DJ 4500 (5)	8,73%	$4,57 \cdot 10^{16}$	$2,35 \cdot 10^8$	26,78%	305,38%	48,51%	36,65%	36,75%
DJ 4500 (6)	55,20%	$8,67 \cdot 10^{22}$	61,95%	62,55%	67,20%	60,04%	56,01%	56,25%

Таблиця 4.20

Похибки прогнозування індексу S&P 500 методами СЛМ та нейронних мереж

	Markov MAPE	1	2	3	4	5	6	7
S&P 2000 (1)	15,35%	71,60%	173177%	678,74%	56,40%	31,85%	28,93%	30,27%
S&P 2000 (2)	12,64%	$3,56 \cdot 10^{29}$	181044%	37,70%	43,95%	50,45%	21,27%	22,46%
S&P 2000 (3)	15,43%	23,44%	98055,74%	50,01%	13297,48%	43,08%	20,19%	22,24%
S&P 2000 (4)	13,21%	33,15%	606,61%	24,76%	$1,5426E+17$	34,36%	16,01%	18,56%
S&P 2000 (5)	54,90%	79,57%	9027,57%	64,04%	22520%	179,03%	51,88%	51,70%
S&P 3000 (1)	11,78%	74,85%	205,37%	59,81%	$5,3 \cdot 10^{31}$	43,08%	26,04%	25,32%
S&P 3000 (2)	14,96%	132,92%	1246,03%	41,80%	63,46%	$3,257E+05$	19,10%	19,12%
S&P 3000 (3)	28,68%	40,20%	28,53%	34,88%	73,58%	33,59%	20,03%	19,62%
S&P 3000 (4)	18,56%	52,33%	431,70%	$1,87 \cdot 10^{54}$	$1,0 \cdot 10^{36}$	74,02%	17,76%	20,12%
S&P 3000 (5)	19,07%	167,00%	31,07%	53,43%	$8,14 \cdot 10^{29}$	37,69%	38,07%	37,67%

Продовження таблиці 4.20

1	2	3	4	5	6	7	8	
S&P 4500 (1)	11,72%	29,90%	23,25%	32,50%	47,94%	63,65%	30,25%	30,31%
S&P 4500 (2)	10,06%	$7,3 \cdot 10^{54}$	19,24%	$8,3 \cdot 10^{96}$	210,98%	41,11%	14,59%	14,59%
S&P 4500 (3)	25,24%	$8,99 \cdot 10^{73}$	41,05%	49,93%	1025,75%	128,06%	34,58%	33,87%
S&P 4500 (4)	25,66%	$4,4 \cdot 10^{119}$	24,66%	40,77%	76,70%	44,20%	11,93%	14,27%
S&P 4500 (5)	22,83%	61,21%	39,04%	57,66%	$9,8 \cdot 10^{71}$	198,86%	29,53%	32,87%

Аналогічні дослідження були проведені з рядом об'ємів газоспоживання в м. Тернопіль. Результати прогнозування часового ряду газоспоживання методом складних ланцюгів Маркова наведено на рис.4.3. Отримані похибки: $MAE = 8.639,7$; $MSE = 1.3126 \cdot 10^8$, $MAPE = 35\%$.

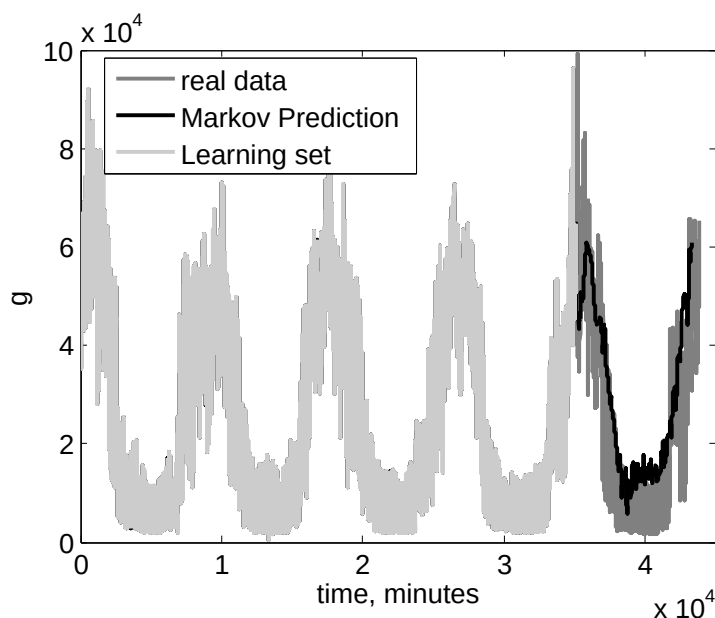


Рис. 4.3. Прогнозування об'ємів газоспоживання у м. Тернопіль.

На основі наведених результатів прогнозування можна зробити наступні висновки: запропонований алгоритм на основі складних ланцюгів Маркова здатний відтворювати характерні риси, які повторюються у реалізації випадкового процесу. Оскільки алгоритм використовує тільки інформацію, наведену в навчальній вибірці то прогнозування, то прогнозувати непередбачених, не наведених у навчальній вибірці рис ряду є неможливим. Можливим шляхом подолання цієї проблеми було б збільшення виборки навчання, проте це спричинює значне збільшення часу, необхідного для розрахунку прогнозу.

4.3. Багатосценарний підхід до прогнозування рядів фінансово-економічної динаміки

У даному розділі представлені результати прогнозування світових фондових індексів, які доступні на сайті <http://finance.yahoo.com>. Експерименти показали, що при різній довжині навчальної вибірки, прогнозні криві дещо відрізняються (див. рис.4.4, 4.5, 4.6 та 4.7). Ряди прогнозів для індексу Dow Jones Industrial Average (далі DJIA) зображено на рис.4.4 є більш корельовані між собою, ніж прогнози індексу FTSE 100, зображеного на рис. 4.6. На рисунках 4.5 та 4.7 вищезазначені графіки усереднено та відкладено одне середньоквадратичне відхилення вверх та вниз від середнього значення. Дані криві можуть інтерпретуватись як довірчі інтервали прогнозу. Час початку прогнозування зображено точкою 1000 та відповідає даті 1 грудня 2011 року.

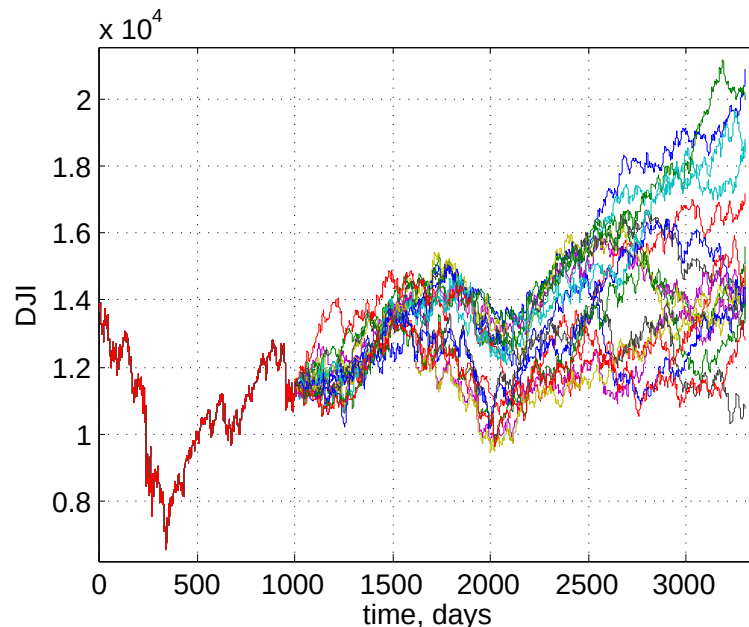


Рис. 4.4. Dow Jones Industrial Average - DJIA (США). Ряд прогнозів, обчислених з різною довжиною навчальної вибірки

Для порівняння індексів різних країн та зображення таких прогнозів на одному графіку пропонується процедура нормалізації. Нормалізовані значення ряду будуть знаходитись у інтервалі $[0...1]$ та обчислюються за наступною формулою:

$$y_n(t) = \frac{y(t) - \min(y(t))}{\max(y(t)) - \min(y(t))}. \quad (4.11)$$

Нормалізовані прогнозні ряди показані на рис. 4.8 (Американські ринки), рис. 4.9 (Європа, зліва ринки розвинутих країн, справа: ринки країн PIIGS), рис. 4.10 (азіацькі

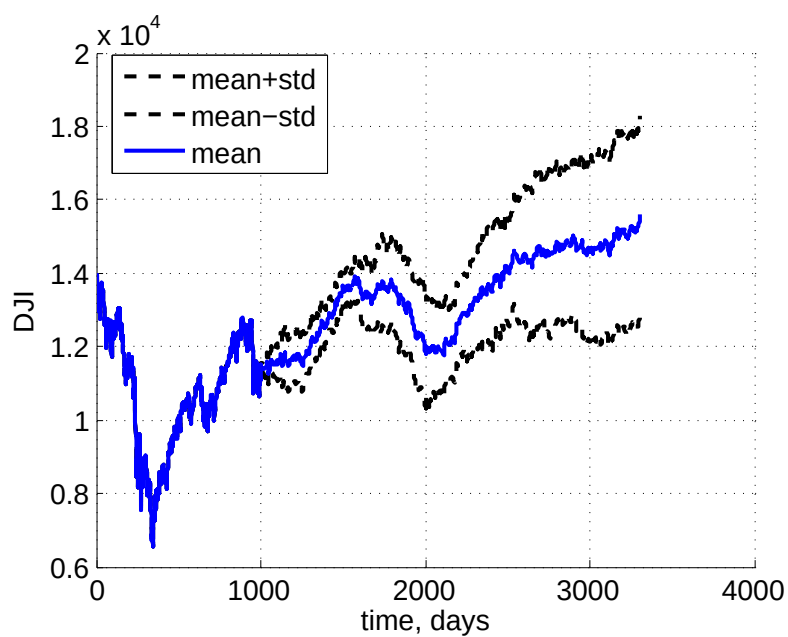


Рис. 4.5. Dow Jones Industrial Average - DJIA (США). Середнє значення та стандартне відхилення прогнозних кривих.

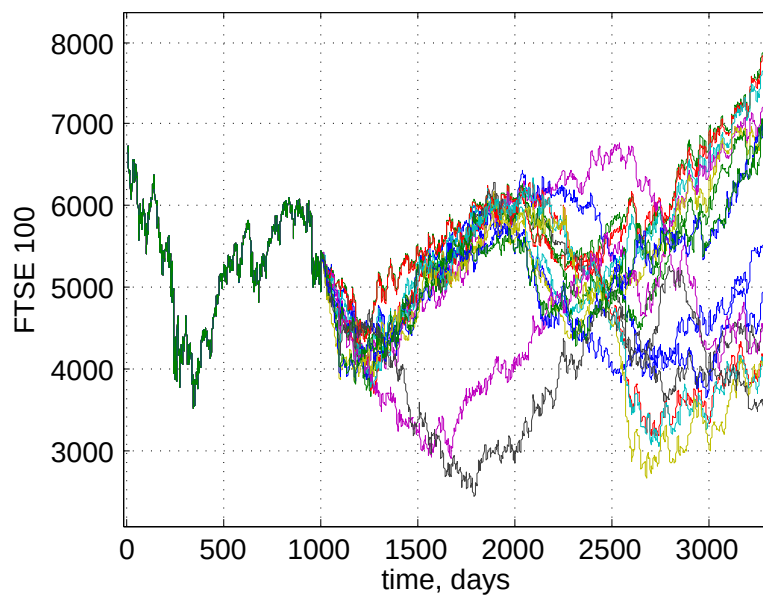


Рис. 4.6. Прогноз індексу FTSE 100 (Великобританія). Ряд прогнозів, обчислених з різною довжиною навчальної вибірки

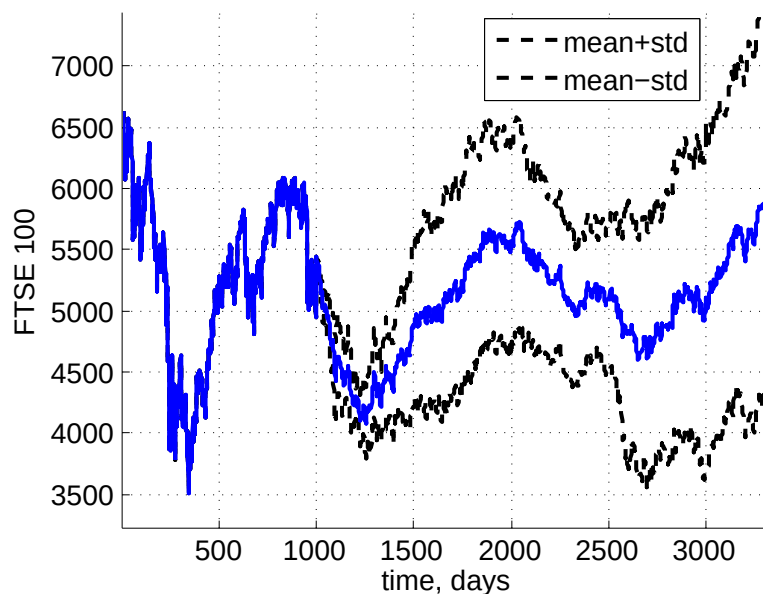


Рис. 4.7. Прогноз індексу FTSE 100 (Великобританія). Середнє значення та стандартне відхилення прогнозних кривих

країни). Усі рисунки містять ряд усередненого прогнозу, які зображені чорною лінією та є середньозваженими з вагами ВВП відповідних країн.

4.4. Обґрунтування та експериментальне визначення параметрів методу дискретного Фур'є-продовження

Приклад апроксимації логарифму індексу Доу-Джонса наведено на рис. 4.12. З рисунку видно досить висока точність апроксимації, така апроксимація дозволяє виокремити низькочастотну складову ряду, що прогнозується, та отримати стаціонарні з точки зору статистики залишки (рис. 4.13) Отримані залишки піддаються аналогічній апроксимації, в результаті чого виділяється наступна низькочастотна складова, наявна в ряді (рис. 4.14). Таким чином, виконуючи деяку кількість ітерацій даної процедури, отримаємо частотний спектр, який містить виокремлені на кожній ітерації низькочастотні складові.

Для прогнозування аналітичну суму гармонічних функцій екстраполюють в часі.

На рис. 4.25 та рис. 4.26 зображено результати прогнозування індексів Доу-Джонса та S&P 500 відповідно, використовуючи різні початкові моменти прогнозу: за 100, 250, 500, 1000, 1500, 2000 днів до кінця відомих даних (останні дані за 12 квітня 2010 року).

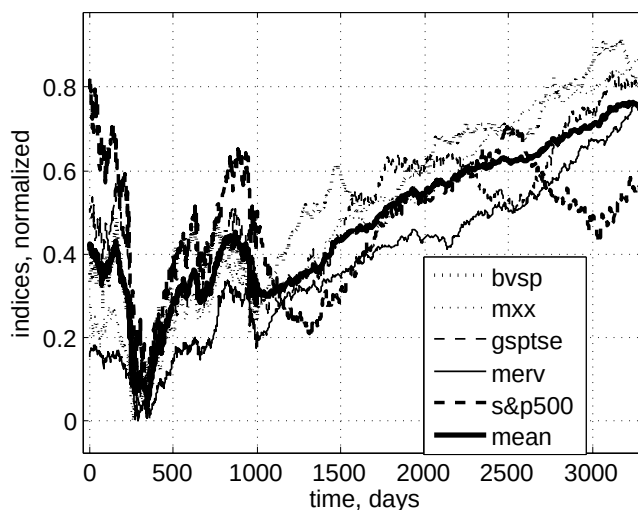


Рис. 4.8. Нормалізовані усереднені прогнози американських фондових індексів. Бразилія (BVSP), Мексика (MXX), Канада (GSPTSE), Аргентина (MERV), США (S&P 500)

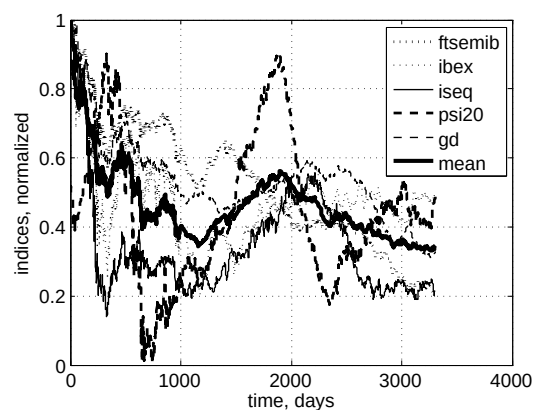
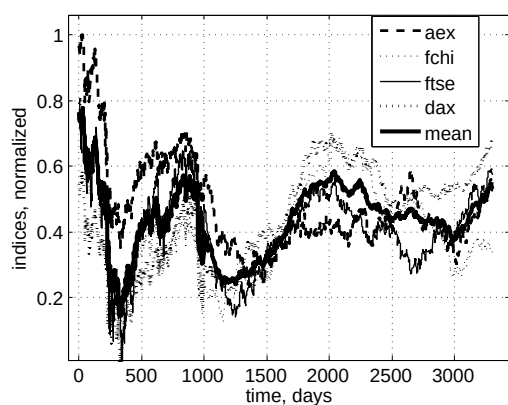


Рис. 4.9. Нормалізовані усереднені прогнози європейських фондових індексів. а) Розвинуті країни: FTSE (Великобританія), DAX (Німеччина) FCHI (Франція), AEX (Нідерланди); б) Португалія (PSI20), Італія (FTSEMIB), Ірландія (ISEQ), Греція (GD) та Іспанія (IBEX)

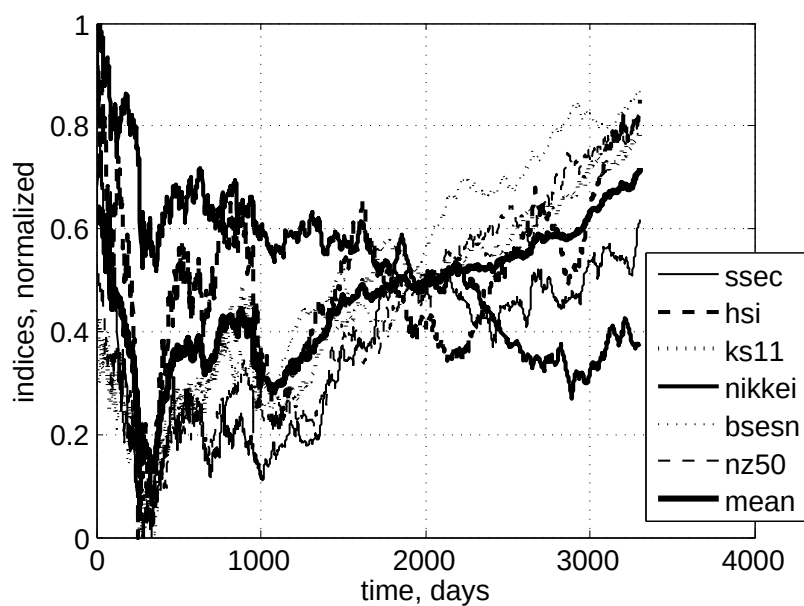


Рис. 4.10. Нормалізовані усереднені прогнози фондових індексів Азії. Китай (SSEC, HSI), Корея (KS11), Японія (NIKKEI), Індія (BSESN), Нова Зеландія (NZ50)

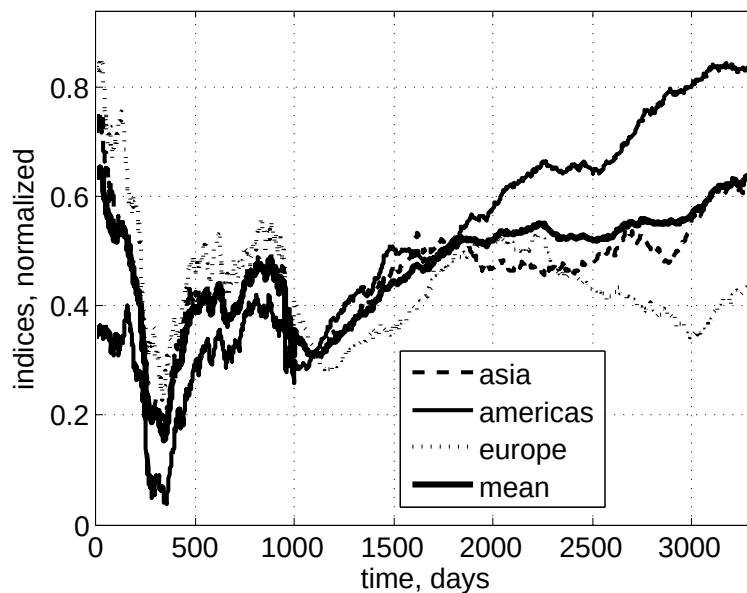


Рис. 4.11. Усереднені та нормалізовані прогнози індексів світових фондових ринків

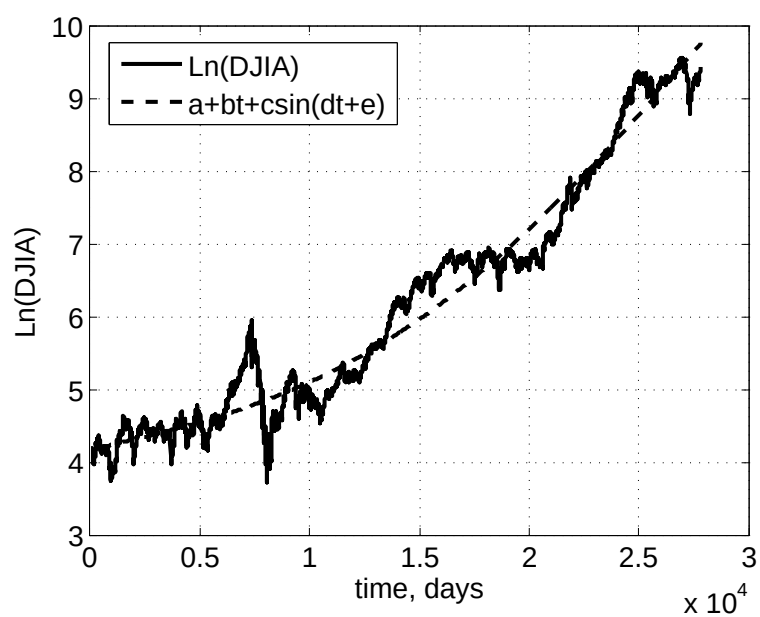


Рис. 4.12. Низькочастотна апроксимація індексу Dow Jones Industrial

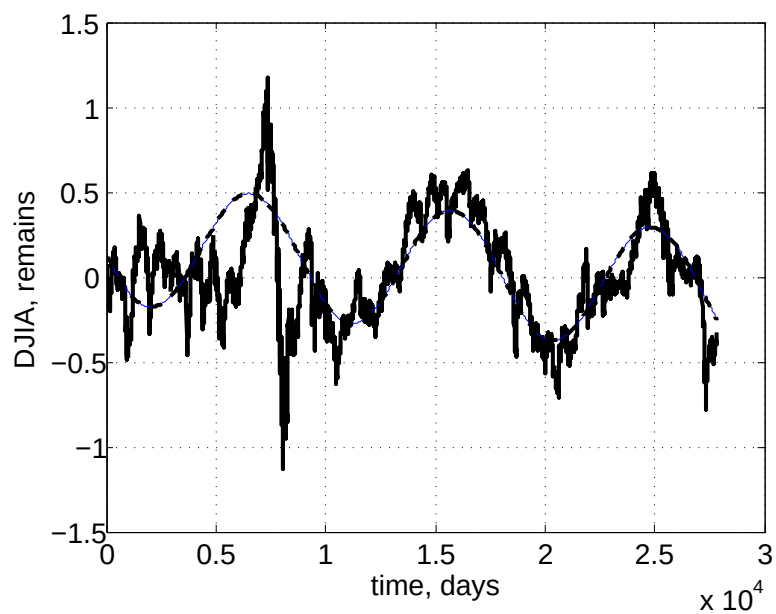


Рис. 4.13. Залишки апроксимації індексу DJIA та їх апроксимація

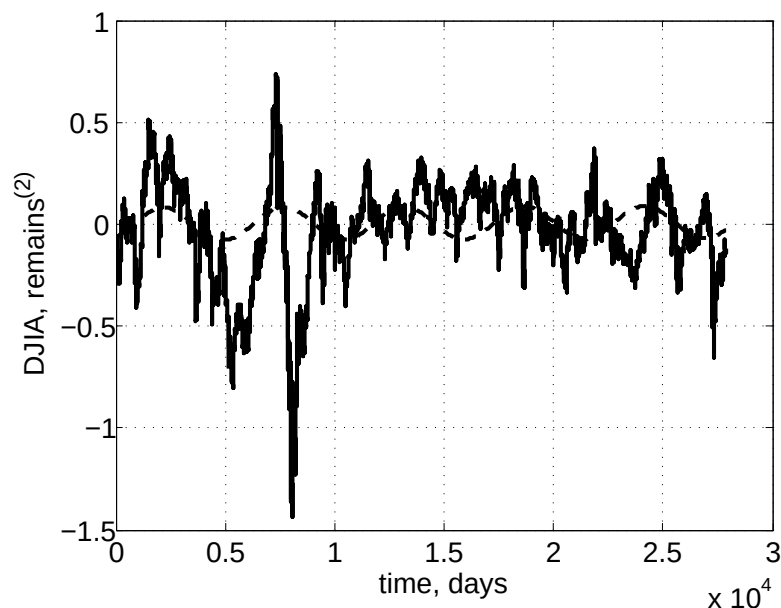


Рис. 4.14. Залишки другого порядку індексу DJIA та їх апроксимація

Дані результати показують достатньо високу точність прогнозів на низькочастотному рівні.

В середньо- та високочастотному діапазоні коливання відбувались нерегулярно, що погіршує результати прогнозів. Середні відхилення вихідного ряду та прогнозів наведені в таблиці 4.22. Критерієм якості прогнозу вибраний середній квадрат похибки MSE (4.4) та MAPE (4.6).

Помічено, що зі збільшенням точності апроксимації збільшується точність прогнозів. Відхилення від поміченої особливості можуть бути пояснені наявністю кризи на навчальному або прогнозному фрагменті, що створює різкі зміни ряду, що досліджується, які погіршують апроксимацію на навчальній вибірці та не дають можливості виявити шукані закономірності. Також кризові явища спричиняють відхилення прогнозів від реального продовження ряду, що може бути пояснене структурними перебудовами в системі, що досліджується. З цього необхідно зробити висновки, що падіння при кризових явищах можуть прогнозуватися наведеною методикою, як видно з рис. 4.25 та рис. 4.26, але глибину падіння визначити не вдається.

Одним із параметрів запропонованого алгоритму, є об'єм навчальної вибірки, який задається кількістю значень вхідного ряду, які використовуються для оцінки параметрів функції (2.7) або (2.8). На рядах, які представлені значною історією значень, таких як індекс Доу-Джонса та S&P 500 проведена експериментальна апробація

запропонованого алгоритму прогнозування з різними довжинами вибірки навчання. Результати експериментальної роботи наведені у таблиці 4.22.

Таблиця 4.22

Дослідження впливу довжини навчальної вибірки на результати прогнозування.

Індекс	Довжина навчальної вибірки	Похибка оцінювання на вибірці навчання		Довжина прогнозу	Похибки прогнозу	
					MSE	MAPE
		MSE	MAPE			
DJ	2000	60,894	3,1 %	1000	525,2870	18%
SP	2000	3,25375	2,27%	1000	33,43603	10,25%
DJ	3000	65,0776	3,5%	1000	1520,364	37%
SP	3000	4,53145	2,97%	1000	50,31766	22,3%
SP	4500	5,53694	3,5%	1000	67,47576	14,5%
DJ	4500	56,7211	3,93%	1000	1170,949	29%
SP	7000	7,10203	4,3 %	1000	284,5858	35%
DJ	7000	63,4337	4,8 %	1000	1533,841	27,4%
SP	10500	11,3889	5,5 %	1000	289,8961	28%
DJ	10500	71,2804	6,1 %	1000	584,3527	12%

Таблиця 1 містить усереднені результати прогнозування з різними довжинами навчальної вибірки (по 5 експериментів з кожним індексом та вибраною довжиною навчальної вибірки). Оптимальна довжина вибірки навчання для прогнозу в 1000 точок коливається в межах 2000-3000. Даний результат був підтверджений на інших часових рядах, але немає можливості провести на них детальне експериментування через недостатність історичних даних. При довжинах вибірки навчання 3000 отримані приклади прогнозів з MAPE=10,25%, що є дещо вищим за результати при інших довжинах вибірки. Це може бути пояснене наявністю постійної частоти коливань, яка залишилась незмінною на прогнозному фрагменті.

4.5. Дискретне Фур'є-продовження з двочастинним наближенням

Метод дискретного Фур'є-продовження був запропонований у роботі [6] та показав достатньо високі прогностичні можливості на спокійних фінансово-економічних часових рядах, що не містили кризових явищ. Були виявлені погіршення ефективності методики при наявності кризових явищ на навчальній чи прогностній вибірці, які спричиняли помилки визначення частоти коливань, нестабільність цих частот та різкі падіння величини, що прогнозується. Такі критичні явища можуть бути спричинені резонансами у коливаннях, наявних у часовому ряді, для виявлення яких необхідно звернути увагу на випадки накладання декількох коливань.

Для покращення прогностичних результатів пропонується включити в апроксимуючу функцію для однієї ітерації дві синусоїди та провести оптимізацію, змінюючи параметри одночасно обох гармонік (двочастинне наближення апроксимуючих функцій). Припускається, що такий підхід дозволить виявити більш складні закономірності, які були пропущені через накладання частот, так звані “биття”.

В таблиці 4.23 наведені параметри функції (2.7) при одно- та двочастинному наближенні, а також значення середнього відхилення MSE.

Порівняння відхилень MSE для двох методів, наведених в таблиці 4.23, показує, що у випадку оцінювання коефіцієнтів одночасно для двох гармонік, мінімальні відхилення отримуються меншими, ніж при повторному застосуванні апроксимації функції (2.7) до залишків. Також видно, що аналітичні функції при одно- та двочастинному наближенні не мають спільних частот, що свідчить про неможливість досягнення точності двочастинного наближення ітераційним застосуванням одночастинного.

На рисунках 4.27 та 4.28 показані графіки апроксимуючих кривих з параметрами, вказаними в таблиці 4.23.

З рисунків видно, що двочастинна апроксимація є більш точною, хоча її частотні складові, взяті окремо, не дають точної апроксимації. Дані результати підтверджують актуальність досліджень, зв'язаних з багаточастинним наближенням гармонічної апроксимуючої функції.

Для прогнозування аналітичну суму гармонічних функцій екстраполюють в часі. Для порівняння прогностичних якостей методів одночастинного та двочастинного наближення, в таблиці 4.24 наведена статистика відхилень прогнозів від реального продовження для різної довжини навчальної вибірки та різних моментів початку

Таблиця 4.23

Результати двочастинного наближення індексів

Індекс	a	b	c_1	d_1	e_1	c_2	d_2	e_2	MSI
S&P2	$2,75 \cdot 10^{-4}$	3,040	0,0741	$1,62 \cdot 10^{-3}$	5,87	0,383	$6,63 \cdot 10^{-4}$	-0,506	105,
S&P2 1 гарм	$2,75 \cdot 10^{-4}$	3,040	0,0741	$1,62 \cdot 10^{-3}$	5,87	0	0	0	184,
S&P2 2 гарм	$2,75 \cdot 10^{-4}$	3,04	0	0	0	0,383	$6,63 \cdot 10^{-4}$	-0,506	112,
S&P1	$2,74 \cdot 10^{-4}$	3,04	-	-	-	0,375	$6,61 \cdot 10^{-4}$	-0,483	113,
S&P1 2ітер	$1,19 \cdot 10^{-6}$	- 0,0066	0,0735	$1,63 \cdot 10^{-3}$	5,78				105,
DJ2	$2,03 \cdot 10^{-4}$	4,65	0,417	$4,10 \cdot 10^{-4}$	1,30	0,657	$5,94 \cdot 10^{-4}$	2,79	622,
DJ2 1 гарм	$2,03 \cdot 10^{-4}$	4,65	0	0	0	0,657	$5,94 \cdot 10^{-4}$	2,79	1960
DJ2 2 гарм	$2,03 \cdot 10^{-4}$	4,65	0,417	$4,10 \cdot 10^{-4}$	1,30	0	0	0	1074
DJ1	$2,33 \cdot 10^{-4}$	4,31	0,402	$6,64 \cdot 10^{-4}$	2,15	-	-	-	1010
DJ1 2і	- $3,13 \cdot 10^{-5}$	1,36	1,18	$9,44 \cdot 10^{-5}$	-2,3	-	-	-	685,

прогнозу. Критерієм якості прогнозу вибраний середній квадрат похибки MSE та MAPE:

Розрахунки прогнозів велись для 10 ітерацій при одночастинному наближенні та для 5 ітерацій при двочастинному. Довжина прогнозів – 1000 денних значень.

Таблиця 4.24 містить усереднені результати прогнозування з різними довжинами навчальної вибірки (по 5 експериментів з кожним індексом та вибраною довжиною навчальної вибірки) для одно- та двочастинного наближення.

Дані результати показують, що у більшості випадків (95%) двочастинне наближення дає кращі прогнози за критерієм середніх квадратів відхилень. Оптимальна довжина вибірки навчання для одночастинного варіанту коливається в межах 2000-3000. При довжинах вибірки навчання 3000 отримані приклади прогнозів з MAPE=10,25%, що є дещо вищим за результати при інших довжинах вибірки. Це може бути пояснене наявністю постійної частоти коливань, яка залишилась незмінною на прогнозному фрагменті.

Таблиця 4.24

Порівняння прогнозів одночастинного та двочастинного наближення

Індекс	Довжина навчальної вибірки	К-ть гармонік в аналіт. ф-ції	Похибка оцінювання на вибірці навчання		Похибки прогнозу	
			MSE	MAPE	MSE	MAPE
DJ	2000	1	60,894	3,1 %	525,287	18%
DJ	2000	2	66,554	3,32%	78,312	2,65%
SP	2000	1	3,254	2,27%	33,436	10,25%
SP	2000	2	3,428	2,38%	9,7438	2,66%
DJ	3000	1	65,078	3,5%	1520,36	37%
DJ	3000	2	63,908	3,56%	96,96	2,9 %
SP	3000	1	4,531	2,97%	50,318	22,3%
SP	3000	2	4,117	2,63%	70,298	22,88 %
DJ	4500	1	56,721	3,93%	1170,949	29%
DJ	4500	2	53,735	3,883%	691,24	24,7%
SP	4500	1	5,537	3,5%	67,476	14,5%
SP	4500	2	4,86	3,22%	96,584	24,99%
DJ	7000	1	63,434	4,8 %	1533,841	27,4%
DJ	7000	2	59,251	4,5%	205,006	5,36%
SP	7000	1	7,102	4,3 %	284,586	35%
SP	7000	2	7,001	4,38%	24,182	4,19%
DJ	10500	1	71,28	6,1 %	584,353	12%
DJ	10500	2	54,558	4,343%	250,468	4,76%
SP	10500	1	11,389	5,5 %	289,896	28%
SP	10500	2	10,43	5,16%	49,838	5,51%

Висновки до розділу 4

В даному розділі наведено експериментальні результати прогнозування часових рядів на основі складних ланцюгів Маркова та дискретного Фур'є-продовження.

Проведено експериментальну апробацію ефективності запропонованого алгоритму прогнозування часових рядів на основі складних ланцюгів Маркова, який базується на дискретному представленні часового ряду та здатний відтворювати характерні риси, які повторюються у реалізації випадкового процесу. Даний метод показав високу, в порівнянні з відомими методами прогнозування, точність прогнозів. В результаті роботи визначено оптимальні параметри методу, такі як кількість дискретних станів s та порядок ланцюга Маркова r , а також розроблено багатосценарний підхід, який дає можливість врахувати невизначеності при прогнозуванні.

Для прогнозування низькочастотної складової запропоновано метод дискретного Фур'є-продовження, який полягає у апроксимації часового ряду на проміжку навчання аналітичною функцією, яка містить тренд та комбінацію гармонічних коливань. Для оцінювання параметрів моделі використовуються нелінійні методи оптимізації для мінімізації функціоналу нев'язки. Експериментальна апробація показала ефективність запропонованого методу для прогнозування рядів фінансово-економічної динаміки. У порівнянні з відомими класичними підходами, дана технологія дає можливість передбачати моменти часу, в яких можлива зміна тенденції часового ряду, що є важливим у роботі трейдерів на фінансових ринках. При високій інформативності отриманих прогнозів, про які говорять відповідність локальних мінімумів часових рядів прогнозу та реального продовження, класичні критерії оцінювання точності прогнозу (MSE, MAPE) не дають можливості виокремити дану перевагу методу у порівнянні з іншими методами прогнозування. Тому актуальним залишається питання розробки нових критеріїв точності прогнозів, які по-перше адекватно відобразили б співпадання точок зміни тренду при прогнозуванні, а по-друге дозволили б спростити задачу пошуку оптимальних параметрів моделі. Похибки прогнозів даного методу можуть бути пояснені наступними моментами: 1) наявністю фінансової кризи на проміжку навчання або прогнозування; 2) похибками моделі частот, що спричиняють невідповідність екстремумів; 3) похибками амплітуд коливань, які не є сталими і можуть збільшуватись під час кризи. Для подолання проблем, які були помічені, пропонується здійснення пошук оптимальних коефіцієнтів одночасно за двома гармоніками (двочастинне наближення) з метою отримати більш складні комбінації гармонічних коливань (накладання частот, "биття"), що покращує ефективність запропонованого методу.

Продовженням роботи має бути удосконалення алгоритму для рядів зі змінними частотами.

ВИСНОВКИ

У дисертаційній роботі вирішена складна науково-технічна задача розробки моделей часових рядів для побудови їх прогнозів з урахуванням невизначеності прогнозів та мультифрактальних властивостей часових рядів та отримані наступні результати.

1. Запропоновано алгоритм прогнозування часових рядів на основі складних ланцюгів Маркова, тобто ланцюгів Маркова з пам'яттю, які, на відміну від відомих моделей часових рядів, враховують ефект пам'яті та здатні виявити закономірності ряду, приховані від класичних методів прогнозування;
2. запропоновано технологію прогнозування згідно ієрархії приростів часу та процедуру “склеювання”, що дозволило відтворювати мультифрактальні властивості часових рядів складних фінансово-економічних систем;
3. удосконалено методи дискретизації часового ряду, враховуючи особливості негаусового розподілу прибутковостей фінансових часових рядів;
4. удосконалено методи апроксимації та екстраполяції часових рядів на основі виокремлення низькочастотного тренду за допомогою дискретного Фур'є-продовження з двохчастотним наближенням;
5. розроблено програмний засіб MarkovChains в системі matlab для прогнозування часових рядів запропонованою методикою;
6. здійснено оптимізацію швидкодії алгоритму за рахунок використання хеш-таблиць;
7. Розроблена web-орієнтована інформаційна система розподіленого прогнозування, що дозволяє суттєво оптимізувати прогнозування ряду показників у реальному часі;
8. виконана експериментальна робота по дослідженню точності отриманих прогнозів і отримано точність $MAPE=22\%$ при середньостроковому прогнозуванні на 1000 торгових днів (приблизно 4 роки).

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

- [1] Yahoo finance. — Електронний ресурс. <http://finance.yahoo.com>.
- [2] КІнто. пфтс. — Електронний ресурс. <http://www.kinto.com/>.
- [3] Soloviev V. Financial time series prediction with the technology of complex markov chains / V. Soloviev, V. Saptsin, D. Chabanenko // Computer Modelling and New Technologies. — 2010. — Vol. 14, № 3. — P. 63–67.
- [4] Чабаненко Д. М. Алгоритм прогнозування фінансових часових рядів на основі складних ланцюгів Маркова / Д. М. Чабаненко // Вісник Черкаського університету. — 2009. — Т. 173, № 173. — С. 90–102.
- [5] Чабаненко Д. М. Виявлення короткочасної та довготривалої пам'яті та прогнозування часових рядів методами складних ланцюгів Маркова / Д. М. Чабаненко // Вісник Національного технічного університету «Харківський політехнічний інститут». Збірник наукових праць. Тематичний випуск: Інформатика і моделювання. — Харків: НТУ ХПІ, 2010. — № 31. — С. 184–190.
- [6] Чабаненко Д. М. Дискретне Фур'є-продовження часових рядів / Д. М. Чабаненко // Системні технології. Регіональний міжвузівський збірник наукових праць. — 2010. — Т. 1 (66), № 1 (66). — С. 114–121.
- [7] Соловьев В. Н. Адаптивная методика прогнозирования на основе сложных цепей Маркова / В. Н. Соловьев, В. М. Сапцин, Д. Н. Чабаненко // Матеріали VI Міжнародної науково-технічної конференції «КОМТЕХБУД 2008» Київ, 21–23 серпня 2008 р. — Київ: Міністерство регіонального розвитку та будівництва України, 2008. — С. 59–60.
- [8] Чабаненко Д. М. Застосування прихованих марківських моделей для дослідження економічних систем / Д. М. Чабаненко, Н. С. Коровяцька // «ПРОБЛЕМИ ТРАНСФОРМАЦІЙНОЇ ЕКОНОМІКИ» Збірник наукових тез II Всеукраїнської науково-практичної конференції. / відповід. ред. к.т.н., доц.. Берідзе Т.М. — Кривий Ріг: КФ ДВНЗ ЗНУ, 2009. — С. 84–85.

- [9] Чабаненко Д. М. Методи нелінійної динаміки для прогнозування економічних систем / Д. М. Чабаненко, В. В. Щерба // «Проблеми економіки: освіта, теорія, практика». Матеріали міжнародної науково-практичної конференції (28 листопада 2008 р.). — Кривий Ріг: Мінерал, 2008. — С. 101–102.
- [10] Soloviev V. Prediction of financial time series with the technology of high-order markov chains / V. Soloviev, V. Saptsin, D. Chabanenko // Working Group on Physics of Socio-economic Systems (AGSOE). — Dresden: 2009. <http://www.dpg-verhandlungen.de/2009/dresden/agsoe.pdf>.
- [11] Soloviev V. Financial crisis phenomena: analysis, simulation and prediction. econophysics approach / V. Soloviev, V. Saptsin, D. Chabanenko // «Humboldt Cosmos: Science and society» (HCS2-Kiev2009, Kiev, 19-22 November, 2009). — Kyiv: 2009. — P. 67.
- [12] Chabanenko D. Discrete Fourier forecasting of economic dynamic's time series / D. Chabanenko, O. Nechinennaya // Information Technologies, Management and Society. Theses of the International Conference. — Riga: 2010. — P. 43–44.
- [13] Soloviev V. Discrete Fourier forecasting of economic dynamic's time series / V. Soloviev, V. Saptsin, D. Chabanenko // Information Technologies, Management and Society. Theses of the International Conference. Information Technologies and Management, 2009 April 16-17. — Riga: Information System Institute, 2009. — P. 43–44.
- [14] Чабаненко Д. М. Дискретне Фур'є-продовження низькочастотних складових рядів економічної динаміки / Д. М. Чабаненко // Системний аналіз та інформаційні технології: матеріали 12-ї Міжнародної науково-технічної конференції SAIT 2010. — Київ: ННК «ІПСА» НТУУ «КПІ», 2010. — С. 336.
- [15] Чабаненко Д. М. Дискретне Фур'є-продовження як алгоритм прогнозування фінансово-економічних часових рядів / Д. М. Чабаненко // Системні дослідження та інформаційні технології № 3, 2012 С. 134–142.
- [16] Чабаненко Д. М. Алгоритми дослідження та прогнозування низькочастотної складової динамічного ряду / Д. М. Чабаненко // Інформаційні технології в освіті, науці і техніці / Матеріали VIII Всеукраїнської конференції молодих науковців ІТОНТ-2010. — Черкаси: ЧДТУ, 2010.

- [17] Чабаненко Д. М. Щодо критеріїв оптимальності в алгоритмах прогнозування фінансових часових рядів / Д. М. Чабаненко // ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ТА МОДЕЛЮВАННЯ В ЕКОНОМІЦІ: Зб. наук. пр. Другої Міжнародної науково-практичної конференції; Черкаси, 19-21 травня 2010 р. / Редкол.: Соловйов В.М. (відп. за випуск) та ін. — Черкаси: Брама-Україна, 2010. — С. 316–317.
- [18] Fama E. F. The behaviour of stock market prices / E. F. Fama // Journal of Business. — 1965. — Vol. 38. — P. 34–105.
- [19] Fama E. F. Efficient capital markets: A review of theory and empirical work / E. F. Fama // Journal of Finance. — 1970. — Vol. 25. — P. 383–417.
- [20] Швагер Д. Технический анализ. Полный курс / Д. Швагер. — Москва: Альпина Бизнес Букс, 2005. — 806 с.
- [21] Мерфи Д. Межрыночный технический анализ / Д. Мерфи. — Москва: Диаграмма, 2002. — 317 с.
- [22] Колби Р. Энциклопедия технических индикаторов / Р. Колби. — Москва: Альпина Бизнес Букс, 2004. — 837 с.
- [23] Удовенко В. А. Форекс. Практика спекуляций на курсах валют / В. А. Удовенко. — Москва: ООО “И.Д. Вильямс”, 2007. — С. 384.
- [24] Грабовецкий Б. Є. Основи економічного прогнозування, Навчальний посібник. / Б. Є. Грабовецкий. — Вінниця: ВФ ТАНГ, 2004. — 162 с.
- [25] Присенко Г. В. Прогнозування соціально-економічних процесів / Г. В. Присенко, Є. І. Равікович. — Київ: КНЕУ, 2005. — 378 с.
- [26] Бідюк П. І. Аналіз часових рядів (навчальний посібник) / П. І. Бідюк, В. Романенко, О. Тимошук. — Київ: Політехніка, 2010. — 317 с.
- [27] Тесля Ю. Н. Введение в информатику природы. Монография. / Ю. Н. Тесля. — Киев: Маклаут, 2010. — С. 255.
- [28] Рассел С. Искусственный интеллект: современный подход / С. Рассел, П. Норвиг; Под ред. . е изд.: Пер. с англ. — Москва: Издательский дом “Вильямс”, 2006. — С. 1408.

- [29] Ежов А. А. Сознание, рефлексия и многоагентные системы / А. А. Ежов // НАУЧНАЯ СЕССИЯ МИФИ–2007. VIII ВСЕРОССИЙСКАЯ НАУЧНО-ТЕХНИЧЕСКАЯ КОНФЕРЕНЦИЯ «НЕЙРОИНФОРМАТИКА–2007»: ЛЕКЦИИ ПО НЕЙРОИНФОРМАТИКЕ. Часть 1. — МИФИ, 2007. — С. 11–51.
- [30] Синергетичні та еконофізичні методи дослідження динамічних та структурних характеристик економічних систем. Монографія / В. Дербенцев, О. Сердюк, В. Соловйов, О. Шарапов. — Черкаси: Брама-Україна, 2010. — С. 287.
- [31] Тобілевич Ю. Агентно-орієнтоване моделювання / Ю. Тобілевич // Вісник Черкаського університету. — 2010. — Т. 173. — С. 79–89.
- [32] Krause A. Evaluating the performance of adapting trading strategies with different memory lengths / A. Krause // ArXiv e-prints. — 2009.
- [33] Лукашин Ю. П. Адаптивные методы прогнозирования временных рядов: учеб. пособ. / Ю. П. Лукашин. — Москва: Финансы и статистика, 2003. — 416 с.
- [34] Єріна А. М. Статистичне моделювання та прогнозування: навч. посібник / А. М. Єріна. — Київ: КНЕУ, 2001. — 170 с.
- [35] Грицюк П. М. Аналіз, моделювання та прогнозування динаміки врожайності озимої пшениці в розрізі областей України : монографія / П. М. Грицюк. — Рівне: НУВГП, 2010. — 350 с.
- [36] Грицюк П. М. До питання про циклічність урожайності зернових / П. М. Грицюк // Моделювання та інформаційні системи в економіці : зб. наук. праць / відп. ред. В.К. Галіцин. — 2008. — Т. 77. — С. 299–314.
- [37] Geisser S. Predictive inference: an introduction / S. Geisser. Monographs on statistics and applied probability. — Chapman & Hall, 1993.
- [38] Голяндина Н. Анализ и прогноз временных рядов: программная реализация метода “гусеница”. — Электронный ресурс. - Режим доступа: <http://www.gistatgroup.com/gus/programs.html>. <http://www.gistatgroup.com/gus/programs.html>.
- [39] Голяндина Н. Метод “Гусеница”–SSA: прогноз временных рядов / Н. Голяндина. — СПб.: ВВМ, 2004. — С. 52.

- [40] Александров Ф. Автоматизация выделения трендовых и периодических составляющих временного ряда в рамках метода “гусеница”-SSA / Ф. Александров, Н. Голяндина // Exponenta Pro. — 2004. — Т. 3,4. — С. 54–61.
- [41] Данилов Д. Л. Главные компоненты временных рядов: метод «Гусеница» / Д. Л. Данилов, А. А. Жиглявский. — СПб: Издательство Санкт-Петербургского университета, 1997. — С. 150. <http://www.gistatgroup.com/gus/book1/manual.html>.
- [42] Черняк О. І. Динамічна економетрика [Текст] : навч. посіб. / О. І. Черняк, А. В. Ставицький. — Київ: КВПЦ, 2000. — С. 120.
- [43] Черняк О. І. Навчально-методичний комплекс з дисципліни “Часові ряди” для студентів спеціальності “Економічна кібернетика” та “Прикладна економіка” / О. І. Черняк, А. В. Ставицький. — Київ: Видавничо-поліграфічний центр “Київський університет”, 2004. — С. 26.
- [44] Цыплаков А. Эконометрический ликбез: прогнозирование временных рядов. / А. Цыплаков // Квантиль. — 2006. — Т. 1. — С. 1–60.
- [45] Бідюк П. І. Моделювання та прогнозування нелінійних динамічних процесів / П. . Бідюк; Ed. by ред. П. І. Бідюк. — Київ: ЕКМО, 2004. — Р. 120.
- [46] Бідюк П. І. Часові ряди: моделювання та прогнозування / П. . Бідюк. — Київ: ЕКМО, 2003. — С. 144.
- [47] Бідюк П. . Методи прогнозування / П. . Бідюк, О. С. Меньяйленко, О. В. Половцев. — Луганськ: Луганський національний ун-т ім. Тараса Шевченка, 2008. — Т. 1. — С. 305.
- [48] Алехин Е. ОСНОВЫ АНАЛИЗА ВРЕМЕННЫХ РЯДОВ. Методические рекомендации студентам факультета экономики и управления ОГУ / Е. Алехин. — Орел, 2005.
- [49] Сапцин В. М. Опыт применения генетически сложных цепей Маркова для нейросетевой технологии прогнозирования / В. М. Сапцин // Вісник Криворізького економічного інституту КНЕУ. — 2009. — Т. 2 (18). — С. 56–66.

- [50] Максишко Н. Моделювання економіки методами дискретної нелінійної динаміки: Монографія / Наук. ред. проф. В.О. Перепелиця. / Н. Максишко. — Запоріжжя: Поліграф, 2009. — С. 416.
- [51] Морозова Т. Г. Прогнозирование и планирование в условиях рынка. Учебное пособие для вузов. Под ред. Т.Г. Морозовой, А.В. Пикулькина. / Т. Г. Морозова, А. В. Пикулькин, В. Ф. Тихонов. — Москва: ЮНИТИ-ДАНА, 1999. — 318 с.
- [52] Присняков В. Ф. Нестационарная макроэкономика: – Учебное пособие. / В. Ф. Присняков. — Донецк: Дон-НУ, 2000. — 209 с.
- [53] Сигел Э. Практическая бизнес-статистика / Э. Сигел. — Москва: Издательский дом «Вильямс», 2002. — 1056 с.
- [54] Растринин Л. А. Экстраполяционные методы проектирования и управления / Л. А. Растринин, Ю. П. Пономарев. — Москва: Машиностроение, 1986. — 120 с.
- [55] Сизов Ю. С. Анализ портфелей ценных бумаг и управление ими в современной России. — Электронный ресурс. - Режим доступу: http://mirkin.eufn.ru/articles_1.htm\#sizov4.
- [56] Алтунин А. Е. Модели и алгоритмы принятия решений в нечетких условиях / А. Е. Алтунин, М. В. Семухин. — Тюмень: Изд-во ТюмГУ, 2000. — 352 с.
- [57] Асаи К. Прикладные нечеткие системы: Пер. с япон. / К. Асаи, Д. Ватада, С. Иваи. — Москва: Мир, 1993. — 368 с.
- [58] Сергеева Л. Н. Клеточные сети с опосредованным взаимодействием в микроэкономическом моделировании / Л. Н. Сергеева // Искусственный интеллект. — 1999. — Т. 2 (специальный выпуск), № 2 (специальный выпуск). — С. 398–406.
- [59] Сергеева Л. Н. Моделирование поведения экономических систем методами нелинейной динамики (теории хаоса) / Л. Н. Сергеева. — Запорожье: ЗГУ, 2002. — 227 с.
- [60] Clements M. P. Forecasting economic and financial time-series with non-linear models / M. P. Clements, P. H. Franses, N. R. Swanson // Int. journal of forecasting. — 2004. — Vol. 20. — P. 169–183.

- [61] Billings S. A. Dual - orthogonal radial function networks for nonlinear time series prediction / S. A. Billings, X. Hong // Neural Networks. — 1998. — Vol. 11, № 11. — P. 479–493.
- [62] Holland J. The dynamics of searches directed by Genetic Algorithms / J. Holland. — Singapore: Word Scientific, 1988.
- [63] Галушкин А. И. Теория нейронных сетей. Кн. 1: Учеб. Пособие для вузов / А. И. Галушкин. — Москва: ИПРЖР, 2001. — 385 с.
- [64] Головкич В. А. Нейронные сети: обучение, организация и применение. Кн.4: Учеб. пособие для вузов / Общая ред. А.И. Галушкина. / В. А. Головкич. — Москва: ИПРЖР, 2001. — 256 с.
- [65] Розенблат Ф. Принципы нейродинамики: Персептрон и теория механизмов мозга. Пер. с англ. / Ф. Розенблат. — Москва: Мир, 1965.
- [66] Сигеру О. Нейроуправление и его приложения. Пер. с англ. под ред. А.И. Галушкина / О. Сигеру. — Москва: ИПРЖР, 2001. — 321 с.
- [67] Рутковская Д. Нейронные сети, генетические алгоритмы и нечеткие системы / Д. Рутковская, М. Пилиньский, Л. Рутковский; Ed. by П. с польск. И. Д. Рудинского. — Москва: Горячая линия -Телеком, 2006. — Р. 452. http://www.sernam.ru/book_gen.php?id=21.
- [68] Снитюк В. Є. Прогнозування. Моделі. Методи. Алгоритми: Навчальний посібник / В. Є. Снитюк. — Київ: Маклаут, 2008. — С. 364.
- [69] Тимчій М. Побудова архітектури нейронних мереж зворотного розповсюдження за допомогою генетичних алгоритмів / М. Тимчій // Вісник Черкаського університету. — 2010. — Vol. 172. — Р. 94–103.
- [70] Poggio T. Theory of networks for approximation and learning / T. Poggio, F. A. Girosi // A.I. Memo. — 1994. — Т. 1140, № 1140. — С. 63 p.
- [71] Вороновский Г. К. Генетические алгоритмы, нейронные сети и проблемы виртуальной реальности / Г. К. Вороновский. — Харьков: Основа, 1997. — 112 с.
- [72] Барский А. Б. Обучение нейросети методом трассировки / А. Б. Барский // Труды VIII Всероссийской конференции «Нейрокомпьютеры и их применение»: Сб. докл. — Москва: Век книги, 2002. — С. 862–898.

- [73] Белим С. В. Математическое моделирование квантового нейрона / С. В. Белим // Труды VIII Всероссийской конференции «Нейрокомпьютеры и их применение»: Сб. докл. — Москва: Век книги, 2002. — С. 899–903.
- [74] Бутенко А. А. Обучение нейронной сети при помощи алгоритма фильтра Калмана / А. А. Бутенко // Труды VIII Всероссийской конференции «Нейрокомпьютеры и их применение»: Сб. докл. — Москва: Век книги, 2002. — С. 1120–1125.
- [75] Гусак А. Н. Подход к послойному обучению нейронной сети прямого распространения / А. Н. Гусак // Труды VIII Всероссийской конференции «Нейрокомпьютеры и их применение»: Сб. докл. — Москва: Век книги, 2002. — С. 931–933.
- [76] Шибхузов З. М. Конструктивный tower алгоритм для обучения нейронных сетей из тп – нейронов / З. М. Шибхузов // Труды VIII Всероссийской конференции «Нейрокомпьютеры и их применение»: Сб. докл. — Москва: Век книги, 2002. — С. 1066–1072.
- [77] Тарасенко Р. О. Метод аналізу і підвищення якості навчальних вибірок нейронних мереж для прогнозування часових рядів: Автореф. дис. канд. техн. наук : 05.13.06 / Одес. нац. політехн. ун-т. — Одеса, 2002. — С. 19.
- [78] Тихонов Э. Е. Совершенствование методов прогнозирования с использованием нейронных сетей и системы остаточных классов: Дисс. к.т.н. / СГУ. — Ставрополь, 2004.
- [79] Широков Р. В. Нейросетевые модели систем автоматического регулирования промышленных объектов: Дисс. к.т.н. / СГУ. — Ставрополь, 2003.
- [80] (ehips) ики ран. генетические алгоритмы. — Электронный ресурс. <http://www.iki.rssi.ru/ehips/Genetic.htm>.
- [81] Исаев С. А. Генетические алгоритмы и машинное обучение. — Электронный ресурс. Режим доступа: http://www.blind.alfint.ru/modules.php?name=News&new_topic=3.
- [82] Лекции по нейронным сетям и генетическим алгоритмам. — Электронный ресурс. <http://infoart.baku.az/inews/30000007.htm>.

- [83] Батищев Д. И. Генетические алгоритмы решения экстремальных задач / Д. И. Батищев. — Воронеж: ВГУ, 1994. — 135 с.
- [84] Ивахненко О. Г. Метод группового учета аргументов - конкурент методу стохастичної апроксимації / О. Г. Ивахненко // Автоматика. — 1968. — Т. 3, № 3. — С. 58–72.
- [85] Ивахненко А. Г. Помехоустойчивость моделирования / А. Г. Ивахненко, В. С. Степашко. — Киев: «Наук. думка», 1985. — 216 с.
- [86] Ивахненко А. Г. Моделирование сложных систем по экспериментальным данным / А. Г. Ивахненко, Ю. П. Юрачковский. — Москва: «Радио и связь», 1987. — 120 с.
- [87] Ивахненко А. Долгосрочное прогнозирование и управление сложными системами / А. Ивахненко. — Москва: Техника, 1975. — С. 312.
- [88] Зайченко Ю. П. Нечеткие модели и методы в интеллектуальных системах. Монография / Ю. П. Зайченко. — Киев: Изд Дом «Слово», 2008. — С. 344.
- [89] Зайченко Ю. П. Нечеткий метод индуктивного моделирования в задачах прогнозирования макроэкономических показателей / Ю. П. Зайченко // Системні дослідження та інформаційні технології. — 2003. — Vol. 3. — Р. 25–45.
- [90] Зайченко Ю. П. Синтез і адаптація нечітких прогнозуючих моделей на основі методу самоорганізації / Ю. П. Зайченко, . О. Заєць // Наукові вісті НТУУ “КПІ”. — 2001. — Vol. 3. — Р. 34–41.
- [91] Granger C. W. J. Forecasting Economic Time Series / C. W. J. Granger, P. Newbold. — San Diego: Academic Press, 1986.
- [92] Zhang Y. Prediction of Financial Time Series with Hidden Markov Models: Master of applied science: Simon Fraser university / School of Computer Science. — 2005. — P. 102.
- [93] Rabiner L. R. A tutorial on hidden Markov models and selected applications in speech recognition / L. R. Rabiner // In Proceedings of IEEE. — 1989. — Vol. 77. — P. 257–286.
- [94] Rabiner L. R. An introduction to hidden Markov models / L. R. Rabiner, B. H. Juang // IEEE ASSP Mag. — 1986. — Vol. June. — P. 4–16.

- [95] Shi S. Taking time seriously: Hidden markov experts applied to financial engineering / S. Shi, A. S. Weigend // In Proceedings of the IEEE/IAFE 1997 Conference on Computational Intelligence for Financial Engineering. — 1997. — P. 244–252.
- [96] Романовский В. И. Дискретные цепи Маркова / В. И. Романовский. — Москва: Гостехиздат, 1949. — С. 436.
- [97] Дынкин Е. Марковские процессы / Е. Дынкин. — Москва: Физматгиз, 1963. — Р. 860.
- [98] Гупал Н. Применение байесовской процедуры распознавания для прогнозирования динамики курсов акций / Н. Гупал, Ю. Матусов // Компьютерная математика. — 2009. — Т. 1. — С. 105–110.
- [99] Приймак М. В. Оцінювання матриць переходів періодичних ланцюгів Маркова / М. В. Приймак, С. Ю. Прошин // Електроніка та системи управління. — 2009. — Vol. 3 (21). — Р. 26–33.
- [100] Приймак М. Похибка оцінок матриць переходів періодичного ланцюга Маркова / М. Приймак, О. Віцентій, С. Прошин // Вісник ТНТУ. — 2010. — Т. 15, № 3. — С. 150–159.
- [101] Капранова Л. Г. Прогноз экспорту української металопродукції методом аналітичного вирівнювання за гармоніками ряду Фур'є / Л. Г. Капранова // Проблемы развития внешнеэкономических связей и привлечение иностранных инвестиций: региональный аспект. — 2010. — Vol. 1. — Р. 240–243.
- [102] Анушина Е. С. Система краткосрочного прогнозирования электрической нагрузки: Ph.D. thesis / Работа выполнена в Санкт-Петербургском государственном электротехническом университете “ЛЭТИ” им. В.И. Ульянова (Ленина). — Санкт-Петербург, 2009.
- [103] Грицюк П. Реконструкція математичної моделі динаміки врожайності за даними часових рядів / П. Грицюк // Моделювання та інформаційні системи в економіці : зб. наук. праць / відп. ред. В.К. Галіцин. — 2009. — Вип. 79. — с. 160–176.
- [104] Буч Г., Рамбо Дж., Джекобсон А. Язык UML. Руководство пользователя = The Unified Modeling Language user guide. / А. Д. Г. Буч, Дж. Рамбо А. Джекобсон ; 2-е изд. — М., СПб.: ДМК Пресс, Питер, 2004. — 432 с.

**ДОДАТОК 1. РЕЗУЛЬТАТИ ТОЧНОСТІ ПРЕДСТАВЛЕННЯ,
КВАНТИФІКАЦІЇ ТА НАСТУПНОГО ВІДНОВЛЕННЯ
ЧАСОВОГО РЯДУ**

Таблиця 4.25

Квантування та відновлення ряду, приріст 0, МАЕ

s	0	1	2	3	4	5
2	60,04811	63,55099	63,52483	57,8836	53,62744	57,95369
3	48,07205	28,49054	46,11746	57,88115	53,62744	48,06744
4	39,32707	30,69113	34,84399	57,88115	37,44608	40,80956
5	42,01063	20,42735	35,65955	57,88115	36,33637	43,02885
6	48,38131	21,32351	39,64088	57,88115	32,43568	49,26477
7	36,06397	18,50938	38,10338	57,88115	35,04669	36,92357
8	31,22053	19,6723	23,07357	57,88115	34,19287	31,61836
9	31,42845	17,21272	31,95812	57,88115	31,75808	32,39787
10	31,09714	16,3421	13,62624	57,88115	31,85516	31,85516
сер. МАЕ	40,85	26,247	36,283	57,881	38,48	41,324

Таблиця 4.26

Квантування та відновлення ряду, приріст 1, MAE

s	0	1	2	3	4	5
2	28,7123	29,03286	29,04159	45,8535	37,35788	28,7123
3	22,42224	35,34859	19,67651	45,8535	37,35788	22,42251
4	13,62889	21,99524	24,08568	45,85327	24,931	15,71549
5	15,39498	26,67345	27,95582	45,85325	20,02739	18,65748
6	16,50274	32,20992	17,52374	45,85325	24,34137	15,81712
7	19,80699	26,48436	41,15324	45,85278	22,8182	18,55593
8	17,97275	25,82118	21,95639	45,85278	18,29971	16,15618
9	18,67071	13,44835	32,76122	45,85278	22,17759	20,53648
10	22,16353	14,9363	24,15115	45,85278	21,07525	21,07525
сер. MAE	19,475	25,106	26,478	45,853	25,376	19,739

Таблиця 4.27

Квантування та відновлення ряду, приріст 0, MAE

s	0	1	2	3	4	5
2	48,31923	59,11249	53,95231	44,17274	26,9297	33,43345
3	28,2323	12,34433	24,0412	44,14046	26,9297	28,66057
4	21,99067	25,83231	23,9021	44,14046	12,18633	23,78298
5	19,42113	8,169812	14,19924	44,14046	8,544357	20,48138
6	13,36357	16,03527	55,64854	44,14046	9,127673	15,86139
7	13,30016	8,225185	19,83945	44,14046	9,396903	15,24754
8	6,894006	12,64329	59,2735	44,14046	7,050711	16,79432
9	6,158324	7,029617	26,18688	44,14046	7,740759	9,433507
10	5,212219	8,679105	59,22746	44,14046	6,516853	6,516853
Сер. MAE	18,099	17,5635	37,3634	44,144	12,7137	18,9124

Таблиця 4.28

Квантування та відновлення ряду, приріст 1, MAE

s	0	1	2	3	4	5
2	34,01241	33,05235	33,28631	33,14871	27,43608	33,01459
3	18,68972	26,01157	18,68503	33,14871	27,43608	19,27743
4	16,50159	14,47732	21,09489	33,14871	20,00146	16,30567
5	13,66881	20,44389	18,28044	33,14871	15,8716	15,98664
6	13,78484	13,61067	11,97909	33,14871	9,411825	15,39004
7	11,60153	7,432041	30,29211	33,14871	10,99234	9,748042
8	11,99394	9,35433	9,421857	33,14871	8,952766	9,255166
9	9,935422	5,781245	26,83529	33,14871	8,106758	9,483466
10	7,186287	8,399471	11,80595	33,14871	7,919955	7,919955
Сер. MAE	15,2638	15,3959	20,1868	33,1487	15,1254	15,1534

ДОДАТОК 2. РЕЗУЛЬТАТИ ПРОГНОЗУВАННЯ ЧАСОВОГО РЯДУ АRІМА-МОДЕЛЯМИ

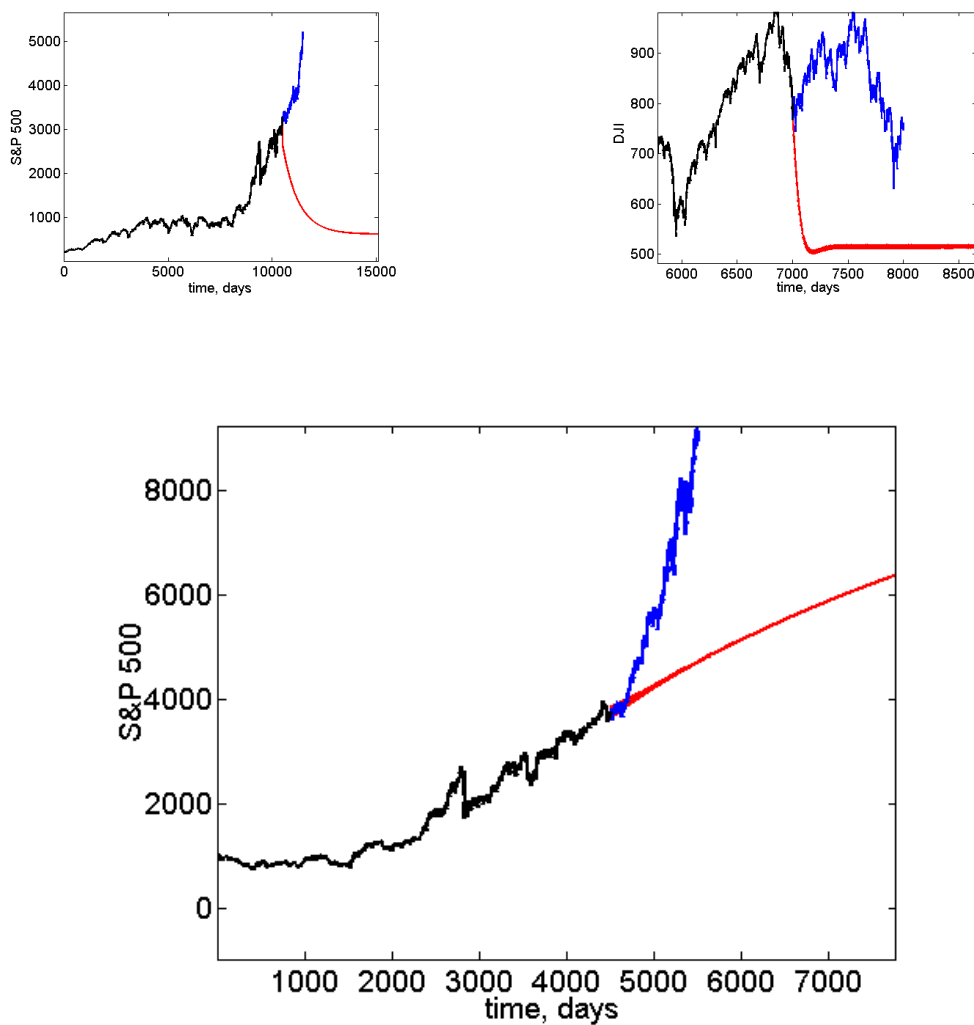


Рис. 4.15. Прогнози індексу Dow Jones Industrial моделлю ARMA(5,1)

ДОДАТОК 3. РЕЗУЛЬТАТИ ПРОГНОЗУВАННЯ ЧАСОВОГО РЯДУ НЕЙРОМЕРЕЖНИМИ МОДЕЛЯМИ

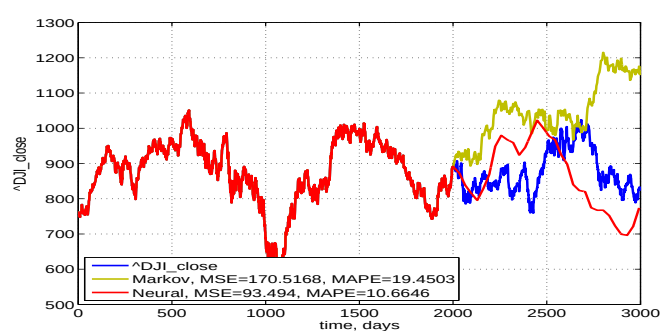


Рис. 4.16. DJIA, learningset: days 10483-12483

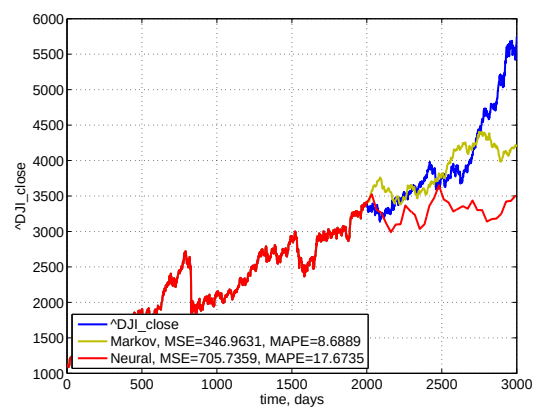


Рис. 4.17. DJIA, learningset: days 13977-12483

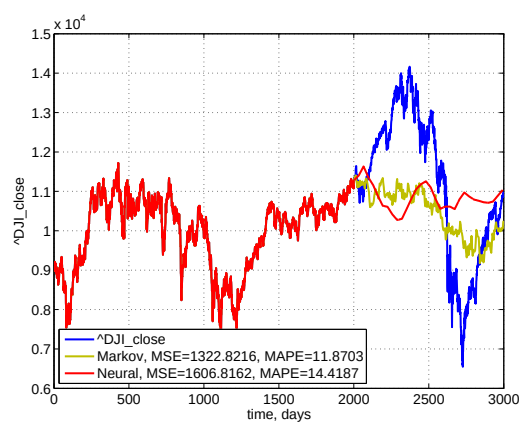


Рис. 4.18. DJIA, learningset: days 17471-19471

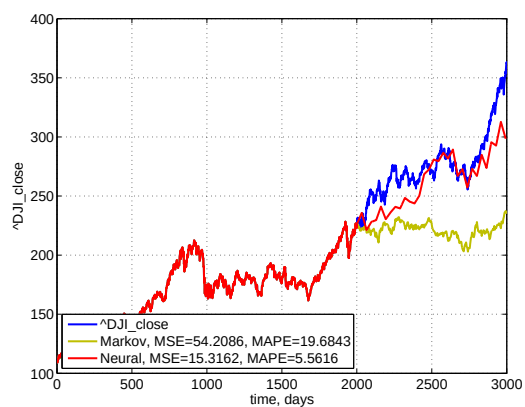


Рис. 4.19. DJIA, learningset: days 3495-5495

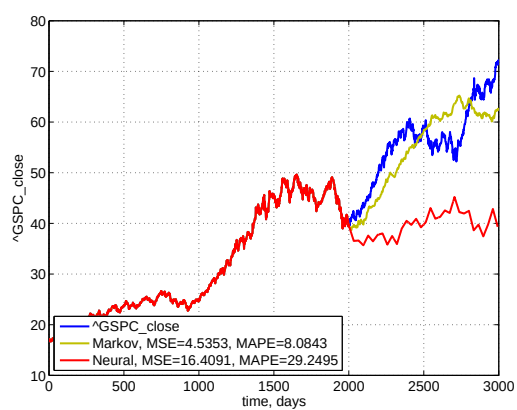


Рис. 4.20. S&P 500, learningset: days 1-2001

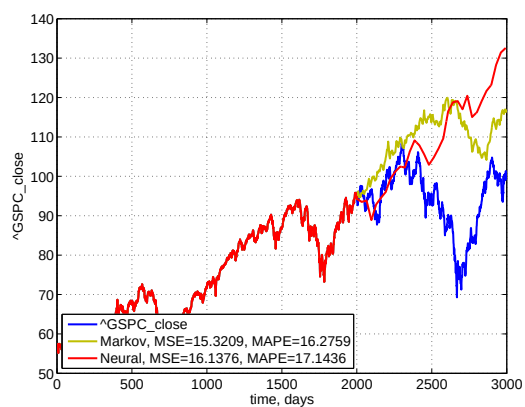


Рис. 4.21. S&P 500, learningset: days 2434-4434

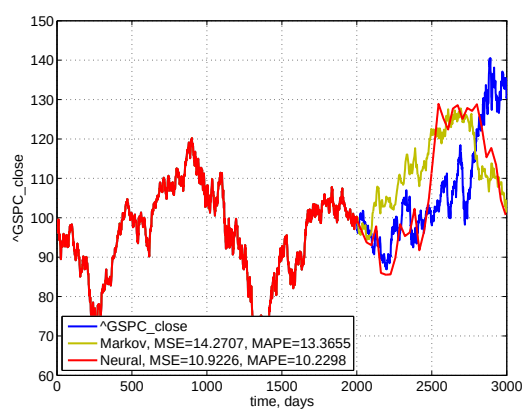


Рис. 4.22. S&P 500, learningset: days 4867-6867

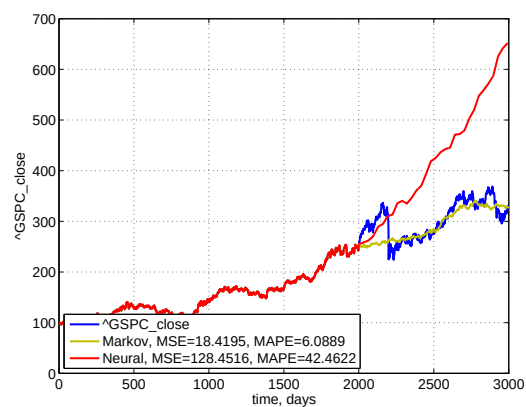


Рис. 4.23. S&P 500, learningset: days 7300-9300

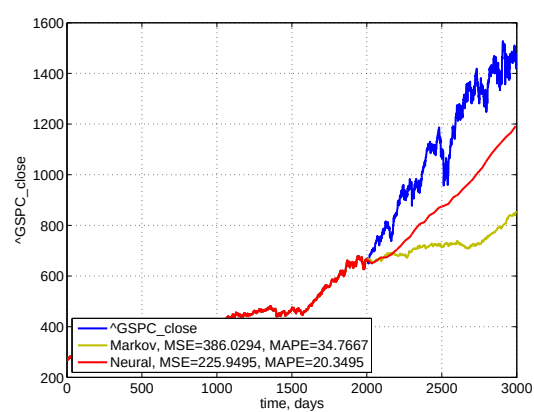


Рис. 4.24. S&P 500, learningset: days 9733-11733

ДОДАТОК 4. РЕЗУЛЬТАТИ ПРОГНОЗУВАННЯ ЧАСОВОГО РЯДУ МЕТОДОМ ДИСКРЕТНОГО ФУР'Є-ПРОДОВЖЕННЯ

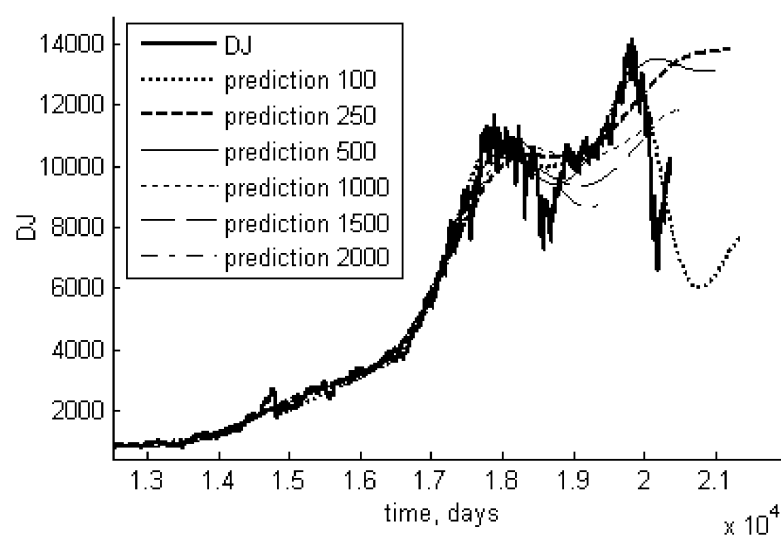


Рис. 4.25. Прогноз індексу Доу-Джонса з використанням дискретного Фур'є-продовження при одночастинній апроксимації

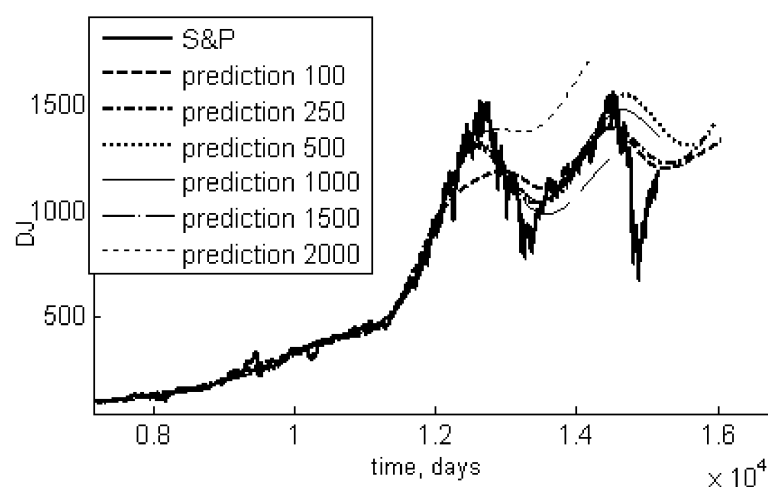


Рис. 4.26. Прогноз індексу S&P 500 з використанням дискретного Фур'є-продовення при одночастинній апроксимації

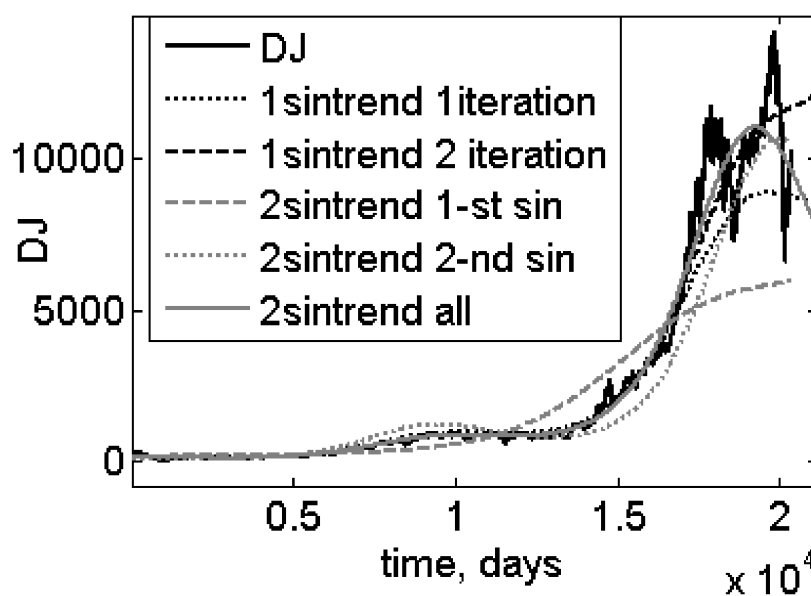


Рис. 4.27. Прогноз індексу Доу-Джонса з використанням дискретного Фур'є-продовення при двочастинній апроксимації

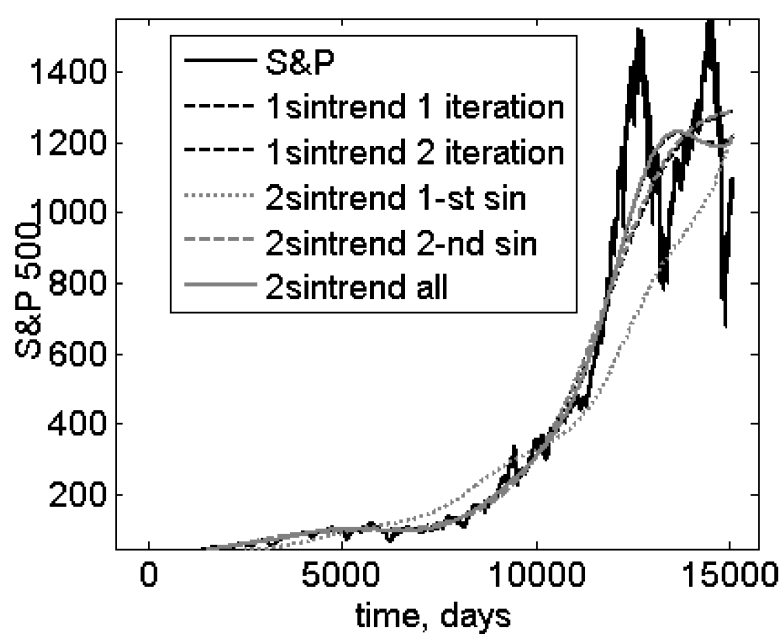


Рис. 4.28. Прогноз індексу S&P 500 з використанням дискретного Фур'є-продовення при двочастинній апроксимації