

Міністерство освіти і науки України
Державний вищий навчальний заклад
«Приазовський державний технічний університет»

О. Ю. Мінц

**МЕТОДОЛОГІЯ МОДЕЛЮВАННЯ ІННОВАЦІЙНИХ
ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ ПРИЙНЯТТЯ РІШЕНЬ В ЕКОНОМІЦІ**

Монографія

(під загальною редакцією чл.-кор. НАНУ, доктора економічних наук,
професора Лисенка Ю.Г.)

*Рекомендовано Вченою радою
ДВНЗ «Приазовський державний технічний університет»*

Маріуполь
2017

УДК 004.89:519.816
М 62

Рекомендовано Вченою радою
ДВНЗ «Приазовський державний технічний університет»
(протокол № 14 від 02.06. 2017 р.)

Рецензенти:

Матвійчук А.В. - д-р економ. наук, проф., професор кафедри економіко-математичного моделювання ДВНЗ «Київський національний економічний університет імені Вадима Гетьмана», директор Інституту моделювання та інформаційних технологій в економіці, генеральний директор Наукового парку КНЕУ.

Жерліцин Д.М. - д-р економ. наук, доц., завідувач кафедру фінансів, банківської справи та страхування ПрАТ ПВНЗ «Запорізький інститут економіки та інформаційних технологій»

Логутова Т. Г. - д-р економ. наук, проф., завідувач кафедри інноватики і управління ДВНЗ «Приазовський державний технічний університет», академік Міжнародної кадрової академії, академік Академії економічних наук України

Мінц О.Ю

М 62 Методологія моделювання інноваційних інтелектуальних систем прийняття рішень в економіці: монографія / О. Ю. Мінц; ДВНЗ «Приазовський державний технічний університет». – Маріуполь, 2017. – 216 с.
ISBN 978-966-604-202-9

У монографії систематизовано та розширено теоретико-методологічні дослідження процесів синтезу інноваційних інтелектуальних систем прийняття рішень. Визначено методологічні підходи застосування методів інтелектуальних обчислень, зокрема нейронних мереж, генетичних алгоритмів, нечіткої логіки для моделювання систем прийняття рішень. Запропоновано моделі і методи вдосконалення процесів прийняття рішень у контексті забезпечення життєздатності та конкурентоспроможності суб'єктів національної економіки.

Для науковців, викладачів, докторантів, аспірантів і студентів економічних спеціальностей вищих навчальних закладів, представників органів державного управління, керівників і менеджерів підприємств та організацій.

УДК 004.89:519.816

ISBN 978-966-604-202-9

© О. Ю.Мінц, 2017
© ДВНЗ «ПДТУ», 2017

ЗМІСТ

ВСТУП.....	4
РОЗДІЛ 1. ТЕОРЕТИКО-МЕТОДОЛОГІЧНІ АСПЕКТИ МОДЕЛЮВАННЯ ІННОВАЦІЙНИХ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ ПРИЙНЯТТЯ РІШЕНЬ.....	8
1.1 Інтелектуальні системи прийняття рішень, як об'єкт дослідження	8
1.2 Методи класифікації задач аналізу і обробки даних в економіці	17
1.3 Методологічні підходи до синтезу інноваційних інтелектуальних систем прийняття рішень	40
1.4 Концепція моделювання інноваційних інтелектуальних систем прийняття рішень в економіці	51
РОЗДІЛ 2. МОДЕЛІ І МЕТОДИ РОЗРОБКИ ІННОВАЦІЙНИХ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ ПРИЙНЯТТЯ РІШЕНЬ.....	62
2.1 Моделювання штучних нейронних мереж для вирішення економічних задач	62
2.2 Нечітко-логічні методи вибору інструментальних засобів інтелектуальних обчислень	82
2.3 Методи оцінки ефективності інтелектуальних методів вирішення економічних задач	95
РОЗДІЛ 3. МОДЕЛІ ІННОВАЦІЙНИХ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ АНАЛІЗУ І ОБРОБКИ ДАНИХ В ЕКОНОМІЧНИХ СИСТЕМАХ	113
3.1 Нейромережеві моделі аналізу даних в економічних системах	113
3.2 Генетичні моделі обробки даних в економічних системах.....	128
3.3 Оцінка параметрів економічних систем методами системно-динамічного імітаційного моделювання	138
РОЗДІЛ 4. РЕАЛІЗАЦІЯ ІННОВАЦІЙНИХ МОДЕЛЕЙ І МЕТОДІВ ПРИЙНЯТТЯ РІШЕНЬ В ЕКОНОМІЧНИХ ЗАДАЧАХ.....	156
4.1 Вибір інструментальних засобів розробки інтелектуальних систем прийняття рішень	156
4.2 Інтелектуальні методи прогнозування показників бюджетного процесу ...	171
4.3 Генетичні методи оптимізації рефлексивних впливів промислових підприємств	181
ЗАГАЛЬНІ ВИСНОВКИ.....	194
ЛІТЕРАТУРА	200
ДОДАТКИ	

ВСТУП

Однією з глобальних сучасних тенденцій розвитку економіки є різке зростання кількості інформації, обсяги якої, за оцінками експертів, підвищуються на 30% щорічно [25]. Лавиноподібне наростання маси різноманітної інформації в сучасному суспільстві отримало назву «інформаційного вибуху», в результаті якого можливості засобів обробки інформації перестали встигати за темпами росту її обсягів. Успіх в інформаційному суспільстві досягається за рахунок найбільш ефективних технологій здобуття, обробки та аналізу інформації. Внаслідок цього на конкурентоспроможність економічних суб'єктів і соціальних груп істотно впливає як нерівність доступу до засобів інформаційних технологій, яка отримала назву «перший цифровий розрив», так і нерівність в знаннях про використання таких технологій – «другий цифровий розрив». Найбільше в даний час значення проблема усунення цифрових розривів набуває в Україні, де відставання в інноваційних технологіях обробки інформації рівнозначно відставанню в розвитку національної економіки.

Розрахунки глобального індексу конкурентоспроможності, які проводяться Світовим економічним форумом, свідчать про погіршення позицій України, яка за підсумками 2015 року займає лише 79 місце в світі серед 140 досліджуваних країн. Питома вага витрат на інновації в загальному обсягу ВВП України складає лише 0.7%, тоді як в розвинених країнах світу становить не менш ніж 2%. Технологічне відставання підтверджується тим, що кількість технологій, переданих за межі України в 3.3 рази менш, ніж придбаних [80].

Разом із тим статистичні дані свідчать про те, що мінімум інноваційної активності пройдено і далі намічається її зростання. За інноваційним підіндексом глобального індексу конкурентоспроможності місце України в 2015 році поліпшилось на 27 позицій. Позитивні тенденції спостерігаються і за показниками інформатизації суспільства, де Україна впритул наблизилася до західноєвропейських показників за кількістю користувачів Інтернет [200].

У цих умовах за необхідне забезпечення ефективного вектору розвитку інноваційних процесів, щоб підготувати якісний стрибок ефективності інновацій. В умовах переходу до інформаційного суспільства особливе значення набуває вдосконалення методів обробки інформації. Світовий досвід свідчить про те, що перспективи цього напрямку полягають в використанні методів інтелектуальних обчислень, які дозволяють відшукувати закономірності в слабоструктурованих даних та набагато підвищують ефективність здобування інформації в умовах «інформаційного вибуху». Зростання обчислювальних потужностей комп'ютерів, появлення нових ефективних алгоритмів призвело до інтеграції інтелектуальних обчислень майже до всіх сфер сучасного суспільства, проте найбільш сильно до цього впливу піддалася економіка. Методи штучного інтелекту є важливою частиною бізнесу таких корпорацій, як Apple, Facebook, Google, Microsoft, які мають ринкову капіталізацію в сотні мільярдів доларів США, та надійно розташовані в верхній частині списку Forbes.

Важливою ознакою сучасних тенденцій розвитку інформаційних технологій є поступове розширення функцій систем на основі інтелектуальних обчислень зі здійснення інформаційної підтримки прийняття рішень людиною до появи повноцінних інтелектуальних систем, здатних до прийняття рішень навіть в умовах мінливого середовища та часткової невизначеності.

Інноваційні інтелектуальні системи прийняття рішень мають самостійне значення для вирішення багатьох економічних задач, зокрема фінансової діагностики, торгівлі на фінансових ринках, створення систем інформаційної безпеки. Але навіть більш значущим слід вважати той факт, що вони здатні стати потужним синергетичним фактором зростання ефективності та життєздатності існуючих економічних систем. В моделі життєздатної системи С. Біра [87] застосування інтелектуальних систем прийняття рішень можливе у всіх п'яти підсистемах, що підвищує їх адаптаційні здібності та життєздатність всієї організаційної структури.

Однак, незважаючи на світові тенденції, перспективи та прогнози розвитку, інноваційні інтелектуальні системи прийняття рішень ще й досі не знайшли широкого розповсюдження в Україні. Серед чинників слід зазначити не тільки інерцію керівників вітчизняних підприємств, які сторожко відносяться до новітніх технологій, які претендують на повноваження в сфері прийняття рішень. Впровадження таких технологій стримує також слабка методологічна база з моделювання інноваційних інтелектуальних систем прийняття рішень, яка б враховувала особливості економічного середовища України, її інформаційний, технологічний та кадровий потенціал.

Таким чином, з урахуванням світових тенденцій, особливостей українського економічного середовища та сучасних завдань, які стають перед економічною наукою слід вважати за актуальну проблему розробку теоретико-методологічних основ і концептуальних положень синтезу інноваційних інтелектуальних систем прийняття рішень, та формалізації їх на рівні моделей і методів. Досягнення цієї мети дозволить вдосконалити процеси прийняття рішень в підприємствах та організаціях України, підвищити ефективність функціонування економічних систем.

В *першому розділі* монографії розглянуто теоретико-методологічні аспекти моделювання інноваційних інтелектуальних систем прийняття рішень. Досліджено генезис інтелектуальних систем прийняття рішень, їх роль і місце в системі управління економічними системами. Систематизовано методи класифікації задач аналізу і обробки даних в економіці. Запропоновано розширену класифікацію таких задач а також інтелектуальних методів їх вирішення. Розглянуто проблеми синтезу інноваційних інтелектуальних систем прийняття рішень та розроблено методологію синтезу із використанням апарату *n*-дольних гіперграфів. Розроблено концепцію моделювання інноваційних інтелектуальних систем прийняття рішень в економіці, що ґрунтується на комплексному використанні інтелектуальних обчислювань на всіх стадіях процесу прийняття рішень

Другий розділ присвячено дослідженню методологічних особливостей моделювання і розробки інноваційних інтелектуальних систем прийняття

рішень. Виділено основні проблеми моделювання штучних нейронних мереж для вирішення економічних задач, серед яких вибір ефективних інструментів моделювання; визначення оптимальної архітектури штучної нейронної мережі; визначення достатнього обсягу та підвищення різноманітності навчальної вибірки; відбір найбільш значущих вхідних даних. Запропоновано шляхи вирішення цих проблем. Розглянуто проблему вибору інструментальних засобів інтелектуальних обчислень та запропоновано використання нечітко-логічних методів моделювання для її вирішення. Досліджено питання оцінки ефективності інтелектуальних методів вирішення економічних задач по відношенню до різних об'єктів і процесів, серед яких алгоритми і програмне забезпечення, процес навчання, а також результати аналізу, або обробки економічних даних.

В *третьому розділі* розглядаються практичні аспекти застосування моделей інноваційних інтелектуальних систем для аналізу і обробки економічних даних. На прикладі вирішення задач розробки фондової торгівельної системи та прогнозування банкрутств українських комерційних банків доводиться залежність економічної ефективності моделювання процесів аналізу даних в економічних системах методами штучних нейронних мереж від постановки завдання. Пропонується використання генетичних алгоритмів для реалізації методу квантування за часом зі змінним кроком для моделювання процесів обробки даних. Виконується оцінка таких параметрів економічних систем, як зовнішні та внутрішні фактори ризику активних операцій комерційних банків методами імітаційного моделювання

Четвертий розділ монографії присвячено різноманітним аспектам реалізації інноваційних моделей і методів прийняття рішень в економічних задачах. Із застосуванням системи математичного моделювання Matlab реалізовано нечітко-логічну модель вибору інструментальних засобів розробки інтелектуальних систем прийняття рішень. Виконано порівняльний аналіз вартості комерційних програмних продуктів для моделювання штучних нейронних мереж та генетичних алгоритмів. Проведено зіставлення програмних реалізацій різних підходів до моделювання генетичних алгоритмів за швидкістю роботи та зручністю використання. Розглянуто використання інтелектуальних методів для прогнозування показників бюджетного процесу. Реалізовано генетичний алгоритм оптимізації рефлексивних управлінських впливів для підвищення конкурентоспроможності комерційних пропозицій промислових підприємств.

Викладені в монографії наукові та методичні положення створюють методологію моделювання інноваційних інтелектуальних систем прийняття рішень в економіці, використання якої дозволяє вдосконалити процеси прийняття рішень господарюючими суб'єктами України та підвищити ефективність функціонування і життєздатність економічних систем за рахунок зменшення невизначеності та вдосконалення їх адаптаційних здібностей.

Монографію підготовлено відповідно до тематики науково-дослідницьких робіт кафедри фінансів і банківської справи ДВНЗ «Приазовський державний технічний університет» за темами Підвищення

ефективності фінансово-кредитного механізму в інноваційному розвитку України (номер держреєстрації 0112U005790, 2012-2013 рр.), Фінансово-кредитне забезпечення стратегії інноваційного розвитку економіки України (номер держреєстрації 0113U007319, 2013-2014 рр.), Підвищення ефективності фінансового управління в умовах нестабільності розвитку національної економіки (номер держреєстрації 0114U004904, 2014-2015 рр.), Удосконалення фінансового управління в Україні (номер держреєстрації 0115U004945, 2015-2016 рр.), а також науково-дослідницьких робіт ВНЗ Укоопспілки «Полтавський університет економіки і торгівлі» за темою «Методологія побудови інноваційних інтелектуальних життєздатних систем управління» (номер держреєстрації 0117U004078, 2016-2017 рр.).

Автор висловлює щирі подяку науковому консультанту, члену-кореспонденту НАНУ, д.е.н. професору Лисенку Ю.Г., рецензентам, д.е.н. професору Матвійчуку А.В., д.е.н. професору Логутовій Т. Г., д.е.н. доценту Жерліцину Д.М., колективу кафедри фінансів і банківської справи Приазовського державного технічного університету за щирі і корисні зауваження та побажання з покращення тексту роботи, колегам, які були причетні до появи цієї праці.

РОЗДІЛ 1.

ТЕОРЕТИКО-МЕТОДОЛОГІЧНІ АСПЕКТИ МОДЕЛЮВАННЯ ІННОВАЦІЙНИХ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ ПРИЙНЯТТЯ РІШЕНЬ

1.1 Інтелектуальні системи прийняття рішень, як об'єкт дослідження

Згідно зі статистичними даними Державного комітету статистики України, понад 93% підприємств України використовують у своїй роботі комп'ютерну техніку. З них в 58,1% підприємств комп'ютери об'єднані в локальну обчислювальну мережу, а 97,4% підприємств мають вихід в Інтернет. Приблизно на 35% підприємств забезпечується бездротовий доступ до корпоративної мережі. 25% підприємств для розміщення інформації використовують внутрішні корпоративні сайти (Інтранет). 33% підприємств мають власний сайт в мережі Інтернет [97].

З підприємств, що мають доступ в Інтернет, 86,3% використовують його для отримання банківських і фінансових послуг, приблизно 80% - для отримання і подачі різних форм звітності, майже половина - для отримання інформації про товари і послуги, 40% - для отримання різних адміністративних послуг [97].

Кількість користувачів Інтернет в Україні стає дедалі більше і в великих містах уже наближається до західноєвропейських показників. Згідно з результатами досліджень Factum Group Ukraine, в 2017 році майже 65% українців у віці понад 15 років користувалися Інтернет не рідше 1 разу на місяць. Причому у великих містах їх частка становить 74%, а в сільській місцевості перевищує 50%. З українців у віці від 15 до 29 років користуються інтернетом 97%, тобто майже всі [200].

На рис. 1.1 показана динаміка зміни кількості користувачів Інтернет в Україні за останні 10 років, зіставлена зі зміною їх активності, індикатором якої прийнято кількість запитів до пошукових систем.

Аналіз графіків на рис. 1.1 дозволяє відзначити зростання активності користувачів Інтернет, кожен з яких генерує все більший обсяг пошукових запитів і відповідно отримує більше інформації.

Виявлені тенденції характерні для сучасного суспільства в цілому. З розвитком і поширенням обчислювальної техніки, відбувається різке збільшення її ролі в процесах прийняття рішень від суто обчислювальних задач і забезпечення швидкого доступу до даних для осіб, що приймають рішення (ОПР) до повноважень в самостійному прийнятті рішень в умовах невизначеності. Сучасні інтелектуальні системи прийняття рішень успішно вирішують задачі найширшого діапазону – від керування автомобілем на міських вулицях до прибуткової торгівлі на валютних і фондових ринках.

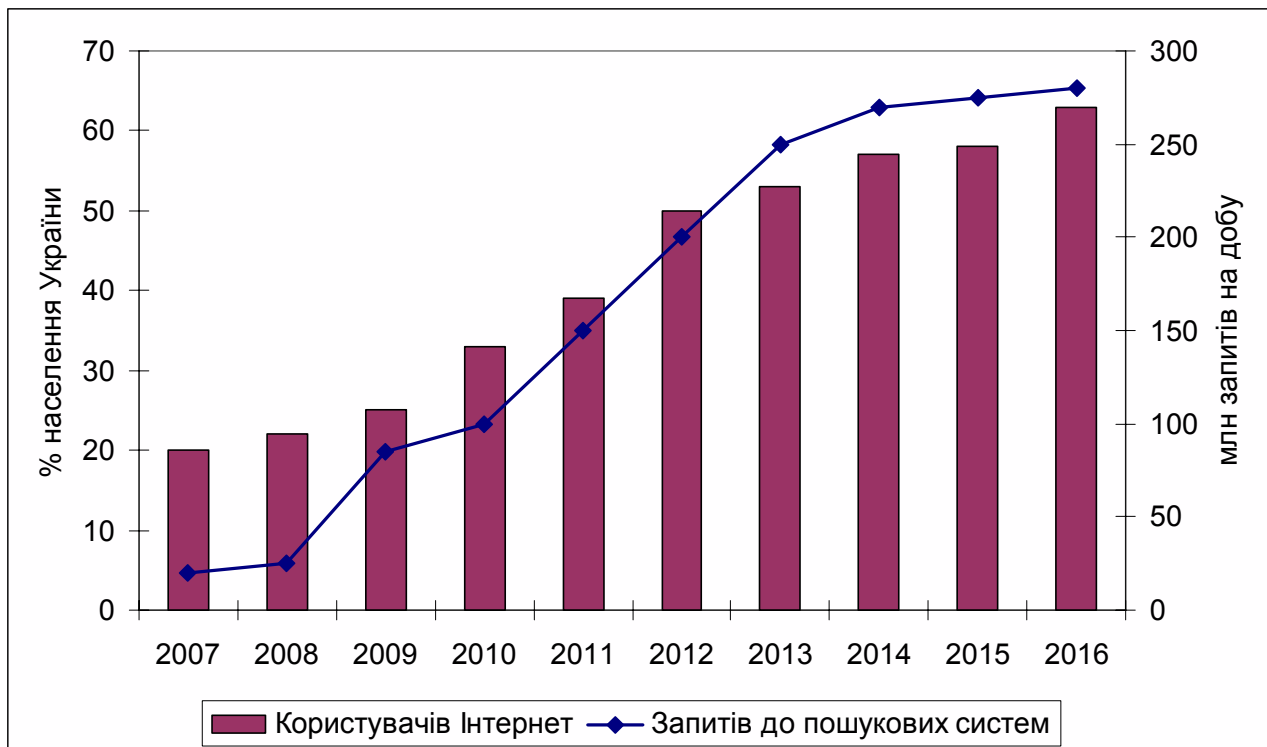


Рисунок 1.1 – Динаміка зміни кількості користувачів Інтернету та їх активності

У процесі прийняття рішень можна виділити кілька етапів. Існуючі в даний час моделі прийняття рішень розрізняються як їх кількістю, так і складом. Наприклад, відома з 1980-х років модель GOFER передбачає п'ять етапів прийняття рішень [41]:

- визначення цілей;
- визначення альтернатив;
- пошук інформації;
- зіставлення позитивних і негативних сторін різних варіантів;
- вибір альтернативи (прийняття рішення) і її реалізація.

Запропонована трохи пізніше модель прийняття рішень DECIDE містить вже шість етапів [21]:

- визначення проблеми;
- визначення обмежень;
- пошук альтернатив;
- прийняття рішення на основі вибору кращої альтернативи;
- розробка заходів по його реалізації;
- моніторинг процесу реалізації рішення і при необхідності його коригування.

У 2007 році П. Браун запропонувала виділяти в процесі прийняття рішення сім етапів [42]:

- загальне визначення мети і очікуваних результатів
- збір даних;
- розробка альтернатив;
- визначення їх переваг та недоліків;
- прийняття рішення;

дії по його реалізації;
пост-аналіз ефективності рішення.

Існують і інші варіанти моделей прийняття рішень, але вони або не мають радикальних відмінностей від перелічених, або, навпаки, носять яскраво виражену специфіку певної предметної області (наприклад, психології).

В узагальненому вигляді у всіх зазначених вище підходах можна виділити три етапи: підготовку прийняття рішення, прийняття рішення, контроль виконання рішення (рис. 1.2). При цьому всі етапи, крім власне прийняття рішення, можуть бути в значній мірі автоматизовані з використанням комп'ютеризованих систем підтримки прийняття рішень (СППР), які були запропоновані і теоретично обґрунтовані в 1960-х роках [130].

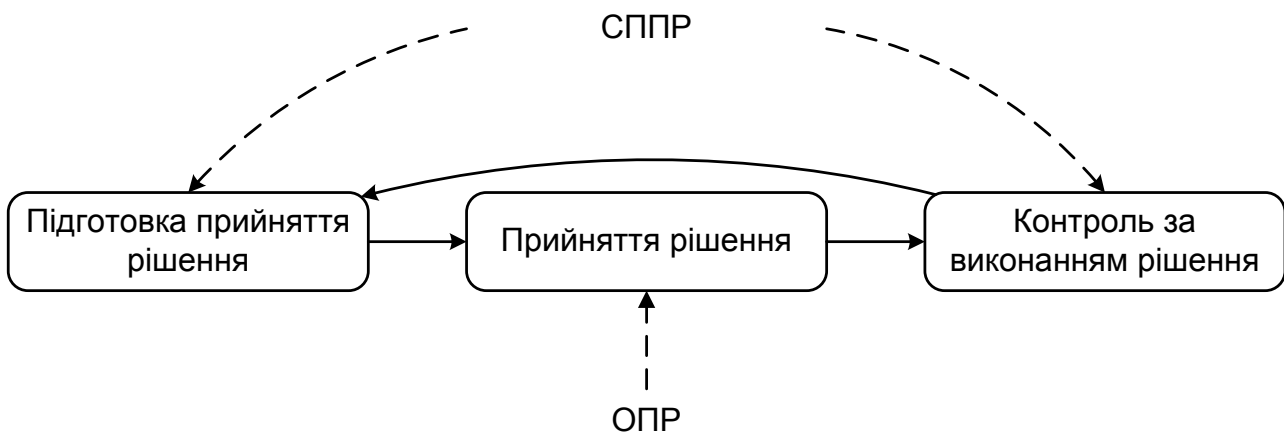


Рисунок 1.2 – Функції ОПР і СППР в процесі прийняття рішень

До цього ж часу відноситься початок масового впровадження на підприємствах автоматизованих систем управління (АСУ) і розвиток експертних систем (ЕС), які орієнтовані на рішення схожих завдань. Тому тривалий час спостерігалася невизначеність використання поняття СППР і сфери застосування таких систем.

В кінці 1960-х років передбачалося, що АСУ зможуть успішно замінити ОПР при вирішенні завдань управління підприємством. Однак досвід практичної реалізації таких систем в 1970-х роках продемонстрував помилковість цього припущення [130]. Аналіз показав, що провали і невдачі в застосуванні АСУ для забезпечення потреб ОПР були обумовлені в першу чергу тим, що на наявному рівні розвитку інформаційно-аналітичних систем було неможливо отримати всю об'єктивну інформацію, необхідну для прийняття рішень. Значна частина взаємозв'язків між об'єктами при цьому спрощувалася, або зовсім не враховувалася.

Використовувані в розглянутих системах моделі прийняття рішень ґрунтувалися на булевій логіці, яка передбачає тільки два можливих стана, наприклад: «так/ні», «відомо/невідомо», «вибрати/відкинути». [116]. У той же час дослідження психології людини і, зокрема, процесу прийняття нею рішень показали, що крім самих фактів, людина оперує і такими категоріями, як ступінь впевненості в них, рівень зв'язку з прийнятим рішенням і іншими.

Можливість урахування цієї інформації з'явилася тільки внаслідок розвитку математичного апарату нечітких множин [70] і нечіткої логіки [113], вперше запропонованих в роботах Л. Заде.

Невдачі автоматизації рішення системних завдань управління зумовили зростання інтересу до СППР в кінці 1970-х років і подальше їх бурхливий розвиток. СППР дозволили об'єднати можливості обчислювальної техніки зі зберігання й обробки інформації із можливостями людини щодо прийняття рішень.

Перші СППР були розроблені для застосування в сфері економіки і фінансів. Вони базувалися на відомих і достатньо надійних економіко-математичних моделях і дозволяли керівникові розраховувати різні варіанти рішень та оцінювати їх результати. Основними компонентами СППР були база даних, система моделей і призначений для користувача інтерфейс. Надалі застосування СППР розвинулося на інші сфери, де для прийняття рішення необхідно проаналізувати альтернативи, порівняти їх і зробити вибір.

Характерними ознаками СППР є робота з неструктурованими та слабо структурованими задачами, відокремлення даних від моделей, інтерактивний режим роботи [20]. Основними завданнями СППР з моменту їх появи є збір і аналіз інформації, з метою зменшення невизначеності в процесі прийняття управлінських рішень.

В середині 1980-х років з'являються комерційно успішні проекти впровадження СППР. Так, в 1987 році компанія Texas Instruments на замовлення United Airlines розробила СППР з управління авіап перевезеннями. Це дозволило значно знизити збитки від польотів і відрегулювати управління різними аеропортами [68].

На подальший розвиток СППР значний вплив мала поява в 2000-х роках нових ефективних методів інтелектуального аналізу інформації та їх інтеграція в технології прийняття рішень, що дозволяє говорити про еволюцію СППР в інтелектуальні системи прийняття рішень (ІСПР), в яких аналіз даних і вироблення рішень здійснюються з використанням інструментарію штучного інтелекту.

Розглянемо класифікацію систем підтримки прийняття рішень.

За способом взаємодії з користувачем виділяють такі СППР [22]:

Пасивні. Забезпечують допомогу ОПР в процесі прийняття рішень, але не можуть забезпечити висунення конкретної пропозиції;

Активні. Безпосередньо беруть участь в розробці правильного рішення;

Кооперативні. Припускають взаємодію СППР та ОПР в процесі пошуку рішення. Пропозицію, яку запропонувала система, користувач може доопрацювати, вдосконалити, а потім відправити назад в систему для перевірки. Після цього пропозиція знову видається користувачеві, і так до тих пір, поки він не схвалить рішення.

ІСПР ґрунтуються на активної, або кооперативної взаємодії.

Класифікація за концепцією СППР побудови, була вперше запропонована в [2] і потім розширена в інших роботах [52, 53]. Згідно з нею виділяють СППР:

Орієнтовані на моделі (*Model-driven DSS*), в основі яких лежать

статистичні, фінансові чи інші моделі, використовуючи які СППР забезпечує перевірку різних варіантів прийнятих рішень. Пік розвитку таких СППР припадає на 1980-ті роки. Надалі більшість функцій модельно-орієнтованих СППР з успіхом стало реалізовуватися за допомогою електронних таблиць типу Microsoft Excel. Проте, в деяких сучасних системах аналізу даних в режимі реального часу (OLAP) можна знайти елементи СППР даної концепції;

Орієнтовані на дані (*Data-driven DSS*), які розраховані, насамперед, на обробку внутрішньої інформації підприємств і організацій, хоча використовуються і зовнішні дані. Інформація може бути представлена як у вигляді часових рядів, так і в іншій формі. Спочатку СППР, орієнтовані на дані, забезпечували швидкий доступ до інформації, її угруповання і складання звітів. Але вже до початку 1990-х років виділяються нові задачі, такі як OLAP, обробка великих масивів даних, *Business intelligence* (методи і інструменти перекладу необробленої інформації в форму, зручну для розуміння і прийняття рішень) та інші. Розвиток СППР цього типу триває й досі;

Засновані на комунікаціях (*Communication-driven DSS*), які використовують інформаційні технології для підвищення ефективності роботи групи користувачів, які займаються вирішенням спільних завдань. Ці СППР в основному забезпечували взаємодії між віддаленими один від одного учасниками проекту (проведення спільних нарад, конференцій, забезпечення внутрішнього листування), проте в даний час цю функцію з успіхом забезпечують інтернет-технології загального призначення. Тому в даний час основним призначенням таких систем є забезпечення спільного доступу і редагування документів, програм, презентацій тощо. Прикладами таких систем є Google Docs, Microsoft Office 365, Apple IWork.com;

Орієнтовані на документи (*Document-driven DSS*), які призначені для роботи з неструктурованою документацією, представленою в самих різних формах, зокрема із сканованими документами, зображеннями, звуковими записами, тощо. Необхідність в таких системах була обумовлена тим, що на момент їх появи лише 5% – 10% ділової інформації було представлено в структурованій формі [17]. Згодом, з розвитком інформаційних технологій, зокрема, технологій розпізнавання образів і глобальних комп'ютерних мереж, частка структурованої інформації істотно зросла, але актуальність СППР, орієнтованих на документи, все ще залишається високою;

Орієнтовані на знання (*Knowledge-driven DSS*), які призначені для формування рекомендацій для прийняття управлінських рішень. При формуванні рекомендацій використовуються бази знань, фактів, правил і процедур. Основу таких систем раніше складали експертні системи, але вже з 1990-х років в них активно впроваджуються методи і інструменти штучного інтелекту.

Аналізуючи наведену класифікацію СППР, Л.Черняк в [235] зазначає тотожність її елементів рівням прийняття рішень в системах Business Intelligence, які виділив Томас Давенпорт (Thomas H. Davenport) в роботі [14]. Ці рівні розташовуються в порядку зростання обсягів оброблюваної інформації і зменшення зв'язності між інформацією та рішеннями. На рис. 1.3 приведена

графічна ілюстрація зв'язку з цим у вигляді піраміди прийняття рішень.

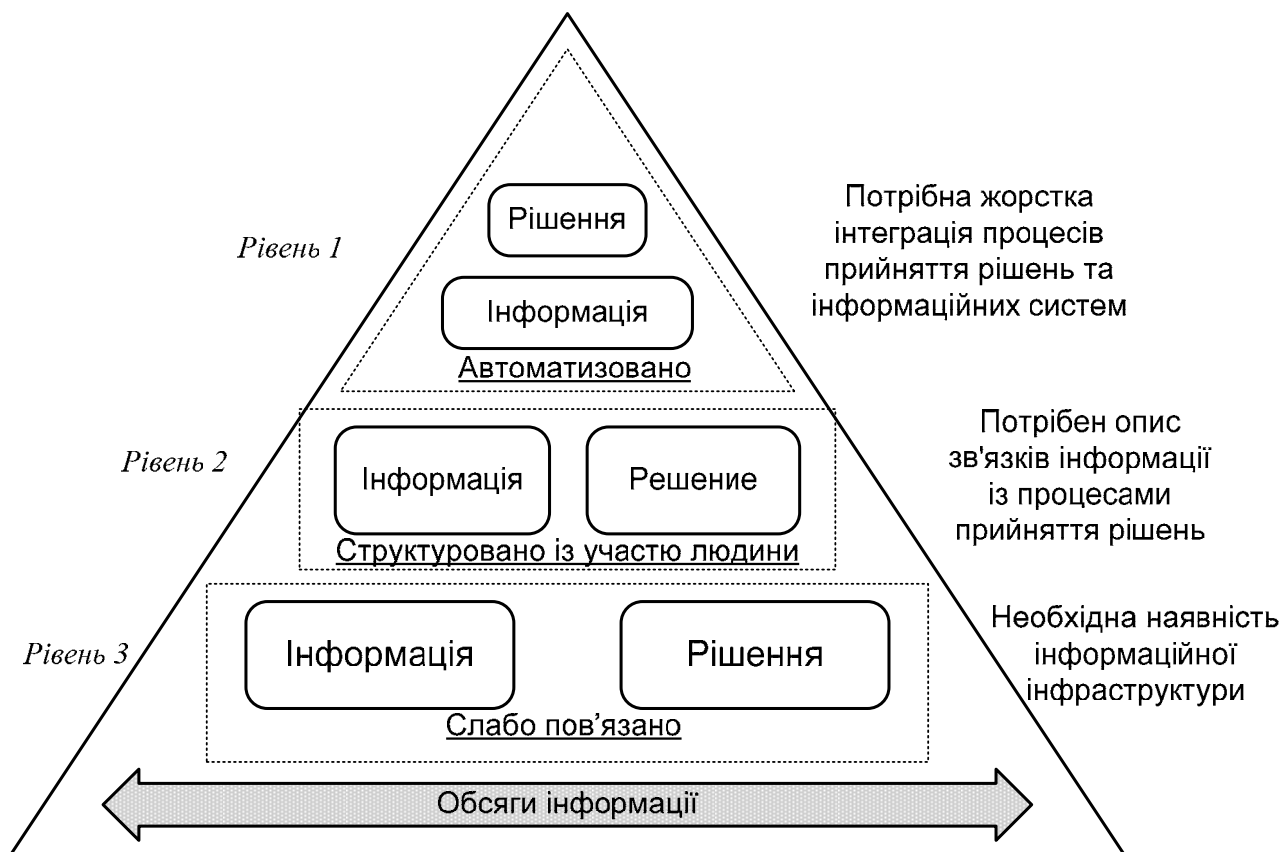


Рисунок 1.3 – Рівні прийняття рішень в системах Business Intelligence

Слід зазначити, що нумерація рівнів піраміди на рис. 1.3. відповідає історичному порядку появи і використання відповідних методів аналізу інформації. Першими стали використовуватися методи, засновані на економіко-математичних моделях, які дозволяють точно описати взаємодію основних чинників, але при цьому не враховують безліч впливів від додаткових факторів. Останніми з'явилися і в даний час активно розвиваються методи аналізу слабоструктурованих даних, які мають малий ступень зв'язку із рішенням, але доступні у великому обсязі, що визначає потребу в комп'ютерах великої обчислювальної потужності для їх аналізу.

В основу цієї класифікації покладено такі принципи [14]:

1. Рішення – результат застосування методів Business Intelligence;
2. Для реального поліпшення якості прийнятих рішень, недостатньо лише доступу до даних і необхідного інструментарію;
3. При підготовці даних для ОПР бажано обмежитися необхідною підмножиною;
4. Відносини між інформацією та рішеннями залежать від типу організації та завдань, що вирішуються. Вони можуть бути слабо пов'язаними, структурованими або автоматичними;
5. Слабка пов'язаність спрощує поставку інформації, але рідко призводить до рішень високої якості;

6. Найбільш перспективним є структурований зв'язок, коли остаточне рішення залишається за людиною, але, при цьому, він діє, використовуючи підготовлені варіанти;

7. Не можна визначити значення Business Intelligence, не зв'язав якимось чином рішення і інформацію, інакше буде незрозуміло, яким інструментом потрібно користуватися;

8. Чим тісніший зв'язок між інформацією та рішенням, тим більш специфічними повинні бути методи прийняття рішень;

9. Спроби пошуку «однієї істини» можуть виявитися занадто дорогими;

10. Якість рішень підвищується при використанні інформаційних технологій, адаптованих до певної області бізнес-аналітики.

Аналіз цих принципів і відповідної їм класифікації систем Business Intelligence, який було проведено в роботі [235] дозволяє безпосередньо зіставити з ними різні концепції побудови систем підтримки прийняття рішень та представити їх у вигляді наступної схеми (рис. 1.4):

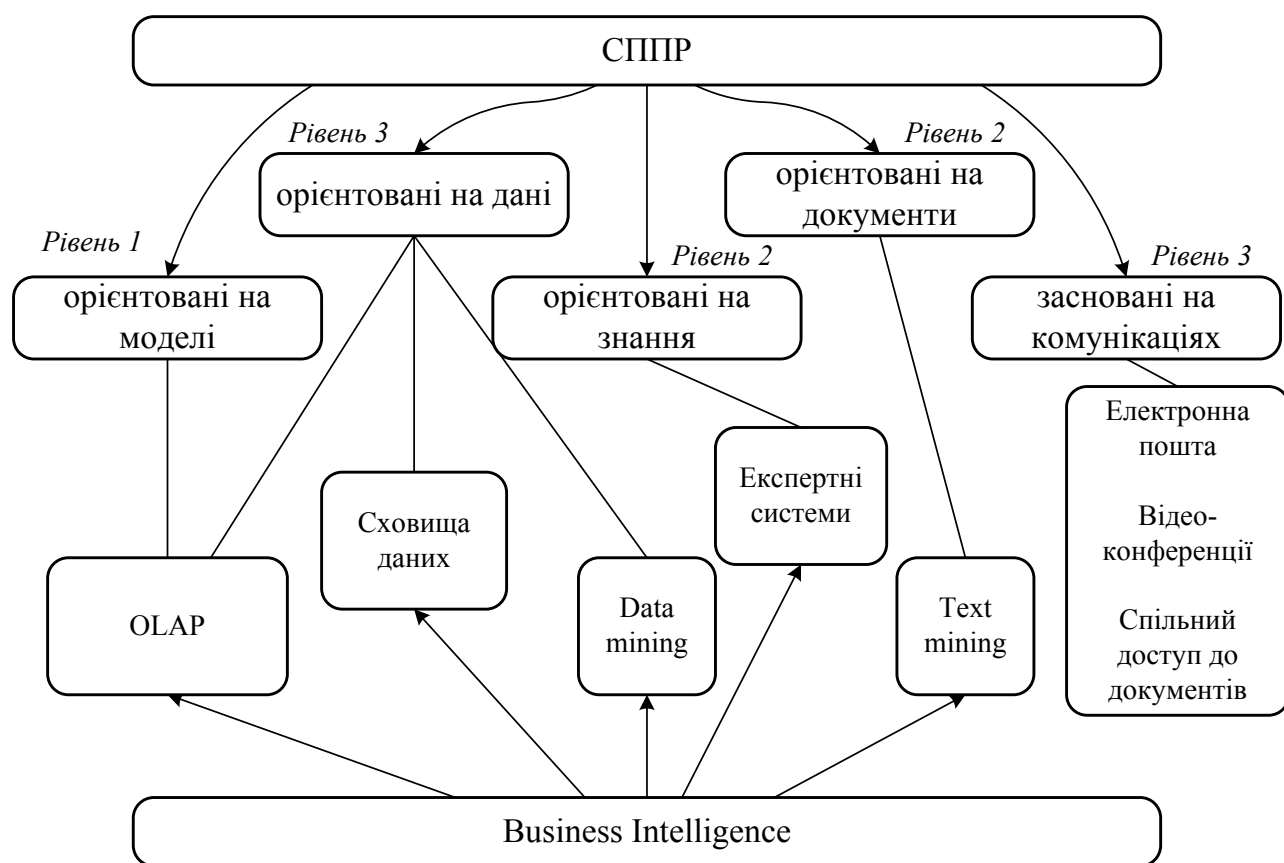


Рисунок 1.4 – Інструментарій Business Intelligence і його відповідність концепціям побудови СППР

Розглянемо схему (рис. 1.4) докладніше.

Верхньому рівню прийняття рішень відповідає модельно-орієнтована концепція побудови СППР. Відповідно до принципів Давенпорта № 8 і № 10, використання готових моделей прийняття рішень можливо тільки в специфічних випадках, але якщо ці умови виконано, модельно-орієнтований

підхід дозволяє отримати найякісніший результат. Основна проблема побудови цих систем полягає в формуванні адекватного набору моделей, які повинні точно описувати бізнес-оточення підприємства, або організації.

Другий рівень систем прийняття рішень і бізнес-аналітики відповідає концепції СППР, що орієнтовані на знання і на документи. Відповідно до принципу Давенпорта № 6, який в свою чергу є наслідком застосування принципів №№ 2,5,7,10, цей рівень прийняття рішень є найбільш збалансованим і перспективним. При цьому з одного боку СППР розвантажує людину, готуючи різні варіанти рішень та враховуючи при цьому велику кількість факторів, а з іншого – цей процес залишається керованим, зв'язки між факторами і ступінь їх впливу на результат прозорі, і остаточне рішення приймає людина.

Третьюму рівню піраміди прийняття рішень (рис. 1.3) відповідають СППР, що орієнтовані на дані та СППР, що засновані на комунікаціях. Відповідно до принципу Давенпорта № 5, на цьому рівні генерується і обробляється найбільшу кількість даних, проте якість рішень, прийнятих на підставі лише слабоструктурованої інформації є досить низьким. Методи аналізу і обробки таких даних в даний час розвиваються найбільш інтенсивно, а сама концепція отримала назву Data Mining (буквально - «Розкопка даних», термін запропонований в 1989 році [201]). Data Mining використовує останні досягнення в сфері штучного інтелекту, машинного навчання та суміжних областей, а обсяги даних, що аналізуються і необхідні обчислювальні ресурси навіть для задач середньої складності виходять далеко за рамки можливостей персональних комп'ютерів. При цьому основною метою аналізу слабоструктурованих даних є представлення їх у вигляді доступному для сприйняття людиною, тобто переклад в структуровану форму, придатну для обробки на рівнях 2 і 1.

Як приклад розглянемо методи прийняття рішень при роботі на валютних і фондових ринках.

Першому рівню прийняття рішень по класифікації Давенпорта відповідає аналіз фундаментальних факторів, тобто таких, для яких можна встановити прозорий причинно-наслідковий взаємозв'язок. Так, програш деякою корпорацією великих судових процесів негативно впливає на курс її акцій, а вихід нового інноваційного продукту, відповідного очікуванням ринку - позитивно, що дозволяє точно змоделювати і автоматизувати процес прийняття рішень. Однак дані, що дозволяють використовувати ці моделі, потрапляють в розпорядження рядових учасників ринку вкрай рідко, що не дозволяє використовувати їх, як основу для біржової діяльності.

Другому рівню прийняття рішень відповідає використання методів технічного аналізу, тобто виділення на графіках руху біржових курсів характерних фігур, що передбачають певні події – падіння, або підйом курсу. Процесом створення таких методів структурування коливань курсу може самостійно займатися кожен учасник ринку. Однак точність прогнозів, зроблених на їх основі, набагато нижче, ніж при фундаментальному аналізі. Крім того, закономірності, які було виявлено, швидко застарівають, у міру того,

як стають відомі іншим учасникам ринку.

Третьому рівню прийняття рішень на валютних і фондових ринках відповідають інтелектуальні системи автоматичної торгівлі, які стали поширені останнім часом. Такі системи використовують «сирі» біржові дані, тобто тільки відомості про коливання ринкових курсів. На підставі цих даних відбувається автоматичне навчання і подальше перенавчання штучних нейронних мереж, які і здійснюють торгівлю на ринку. Такі системи мають ряд переваг перед згаданими вище (можливість цілодобової торгівлі, відсутність «людського фактора» і тому подібні), проте їх недоліком є порівняно низька ефективність. Крім того, складність закладених рішень виключає можливість кваліфікованого втручання в роботу системи з боку користувача.

Очевидно, що всі три розглянутих рівня прийняття рішень не суперечать, а, навпаки, доповнюють один одного, дозволяючи отримати найбільший економічний ефект тільки при спільному використанні.

Таким чином, повноцінна СППР повинна включати всі рівні структурування та обробки інформації - від слабозв'язаного до автоматизованого, при цьому маючи можливість гнучко перемикатися між ними, в залежності від ситуації.

З формальної точки зору задача прийняття рішень є окремим випадком задачі управління і може бути зведена до задачі вироблення управлінських впливів ω_i , на підставі вхідної інформації z_i , яка описує стан системи і її оточення, а також правил, що регламентують прийняття рішень c [100].

Функцію відображення вхідних даних на вихідні позначимо δ :

$$\delta: z_i, c \rightarrow \omega_i. \quad (1.1)$$

Рішення задачі (1.1) на практиці пов'язано з подоланням численних перешкод, одним з яких є складність керованої системи. Поняття «система» в теорії визначається, як сукупність моделей її поведінки Ψ_a и структури Ψ_b , пов'язаних відношенням цілісності $P(\Psi_a, \Psi_b)$:

$$S = \{\Psi_a, \Psi_b, P(\Psi_a, \Psi_b)\}. \quad (1.2)$$

При цьому модель поведінки системи Ψ_a в явному вигляді зазвичай відсутня, що змушує розглядати її у вигляді «чорної скрині», що дає певні відгуки на вхідні сигнали та керуючі впливи. Відгук системи прийнято розглядати як зміна значень показників, що характеризують її стан.

Як відомо з базових принципів теорії управління [81], ідеально керованою є система, яка перебуває в заданому стані з ймовірністю $p=1$. З цього випливає, що для будь-якого стану такої системи має існувати вплив, який гарантовано переводить систему в інший довільний стан. Дане твердження може бути виведено і з принципу різноманітності Ешбі, відповідно до якого різноманітність системи, що управляє, має бути не меншою, ніж різноманітність об'єкта управління [205].

Оскільки стан системи характеризується набором деяких показників, в системі, що ідеально управляється, повинні бути відомі всі можливі набори показників і способи переходу від одного набору до іншого. Якщо система характеризується набором з n показників, кожен з яких може приймати k можливих значень, то загальна кількість станів системи дорівнює k^n . При цьому теоретична межа обчислювальної складності (межа Бремермана), після якого задача стає трансобчислювальною, становить 10^{93} [6], тобто відповідає системі описуваної набором з 93 показників, кожен з яких може приймати 10 різних значень, або системі з 308 показників, кожен з яких може приймати 2 можливих значення. Ця межа також легко досягається при вирішенні NP-повних задач. Так, кількість рішень відомої задачі комівояжера оцінюється формулою $\frac{(n-1)!}{2}$, що робить її принципово нерозв'язною методом прямого перебору вже при $n = 68$ вузлах, оскільки кількість варіантів при цьому стає рівним $1.824 \cdot 10^{94}$.

Рішення таких завдань за практично прийнятний час може бути здійснено тільки за умови штучного обмеження кількості аналізованих варіантів. Таке обмеження може бути зроблено аналітично, на рівнях 1, або 2 піраміди Давенпорта, наприклад, шляхом угруповання і укрупнення системи показників, чи відсікання явно нежиттєздатних варіантів. Однак цей шлях вимагає детального аналізу кожної задачі, із залученням експертів, що збільшує витрати на розробку і супровід СППР і тому не завжди є виправданим. Тому в даний час в сфері бізнес-аналітики і активно розвиваються засоби підтримки прийняття рішень на основі слабоструктурованих даних, які дозволяють в значній мірі формалізувати процедури аналізу. Наприклад, розробка спеціальних методів швидкого пошуку кращого рішення вищезгаданої задачі комівояжера зайняла у математиків більше 150 років. У той же час загальні методи нелінійної оптимізації (наприклад, метод генетичних алгоритмів, або метод імітації відпалювання) дозволяють успішно знаходити досить хороші вирішення цієї задачі (в межах 1% -2% від кращого) на сучасному персональному комп'ютері за кілька десятків хвилин, навіть із урахуванням часу постановки [103].

Таким чином, концепція ІСПР, що базується на використанні сучасних методів аналізу слабоструктурованих даних для прийняття рішень в управлінні складними системами може бути актуальна для вирішення більшості економічних задач, які можуть бути описані формальними ознаками і не засновані на використанні таких суто людських якостей, як інтуїція, харизма і т.п. Використання інтелектуальних методів прийняття рішень дозволяє знизити витрати на розробку і супровід системи управління і підвищити їх точність.

1.2 Методи класифікації задач аналізу і обробки даних в економіці

Відповідно до сучасних концепцій розвитку суспільства, індустріальна епоха не є останньою в розвитку людства. Так, Е. Тоффлер виділяє три хвилі в розвитку суспільства. Перша – аграрна, друга – індустріальна, а третя –

інформаційна, якій відповідає суспільство, що засноване на знанні [221]. Однією з характерних рис інформаційного суспільства, яка була передбачена теоретично і згодом знайшла практичне підтвердження, є превалювання знань над капіталом. Тобто успіх в ньому досягається за рахунок найбільш ефективних технологій генерації та обробки інформації. Яскравим прикладом виникнення таких технологій в економічній сфері є краудфандінг і краудсорсінг, які неможливі без ефективних засобів інформаційного обміну [152]. На конкурентоспроможність економічних суб'єктів і соціальних груп істотно впливає як нерівність доступу до інформаційних технологій, яка отримала назву «перший цифровий розрив» (англ. *digital divide*) [168, 215], так і нерівність в знаннях про використання таких технологій (другий цифровий розрив) [219, 50].

Масове впровадження комп'ютерних технологій призвело до того, що кількість інформації, що потребує обробки, збільшується в геометричній прогресії [25] приблизно на 30% щорічно. Лавиноподібне зростання маси різноманітної інформації в сучасному суспільстві отримало назву «інформаційного вибуху» [223].

Всього 50 років тому для аналізу економіки було достатньо лінійних моделей, які реалізовувалися з використанням обчислювальних засобів порівняно невисокої потужності. У теперішній же час для вирішення деяких економічних задач необхідно використовувати розподілені обчислювальні мережі, що складаються з тисяч високопродуктивних комп'ютерів.

Енциклопедичний словник економіки та права, виданий в 2001 році [239], визначає економічні задачі, як «задачі, які вирішуються в процесі економічного аналізу, планування, проектування пов'язані з визначенням шуканих невідомих величин на основі вихідних даних». При цьому номенклатура економічних завдань, що вирішуються з використанням обчислювальної техніки, за останній час істотно розширилася. Комп'ютерне моделювання починає застосовуватися в таких областях, де раніше покладалися тільки на досвід і інтуїцію людини.

Зазначені фактори обумовлюють необхідність класифікації і систематизації завдань аналізу і обробки даних в економіці. Поняття «систематизація» і «класифікація» раніше часто розглядалися і використовувалися як синоніми. Але в даний час їх все частіше поділяють.

Класифікацією називають об'єднання понять, явищ, предметів в групи на підставі притаманних їм загальних ознак [226].

Систематизацією називають розташування класів, понять, предметів і явищ в певному порядку, відповідному відносин між ними [213].

Відповідно до цих визначень «класифікація» може розглядатися як підпорядковане поняття, по відношенню до «систематизації». Однак це справедливо лише в тому випадку, якщо об'єктами систематизації є класи, а її підставою є природні ознаки. Якщо систематизація проводиться по штучним ознаками (наприклад – за алфавітом), вона не може бути використана для отримання нових знань і її наукова цінність відсутня [213].

Аналіз [8, 17, 23, 195], та інших джерел, присвячених обробці даних, показав, що більшість авторів не приділяють проблемам класифікації та

систематизації достатньо уваги, обмежуючись невеликою кількістю аналізованих ознак. В окремих випадках можна спостерігати зворотну логіку побудови дослідження, де спочатку описуються методи і інструменти, а вже потім задачі, які вирішуються за їх допомогою.

У той же час систематизація та класифікація є найважливішими компонентами наукових досліджень і в ряді джерел розглядаються, як одні з головних функцій науки [226]. Роботи по систематизації знань сприяють розвитку науки і переходу її з емпіричного на системний рівень. Якщо в основу класифікації покладено об'єктивно існуючі ознаки, то вона сама по собі може бути інструментом наукового пізнання і служити для отримання нових знань і закономірностей.

Розглянемо класифікацію задач аналізу і обробки даних. Метою такого дослідження є виділення оптимальних методів вирішення цих задач. Тобто, виходячи з інформації про набір вхідних даних, очікувані результати та інших супутніх умов повинен забезпечуватися вибір оптимального інструменту вирішення. При цьому повинні бути враховані різні критерії оптимальності, до яких можна віднести швидкість побудови системи прийняття рішень, швидкість обробки даних, точність рішення, витрати (загальні і з розбивкою по компонентах), необхідні обчислювальні ресурси і таке інше.

Розглянемо існуючі підходи до класифікації завдань аналізу і обробки даних. При цьому маються на увазі переважно економічні задачі, хоча, завдяки розвитку інформаційних технологій, в даний час подібні задачі можуть виникати в будь-якій сфері.

Підхід, заснований на розгляді аналізу, обробки даних і супутніх завдань, як окремих фаз бізнес-циклу реалізації ІСПР, назовемо процесним. Відповідно до нього ці задачі можна класифікувати в такий спосіб [49, 8]:

- *Аналіз предметної області.* Для опису цього етапу в англійській літературі вживається термін *business understanding*, тобто «розуміння бізнесу». У цій фазі здійснюється загальний аналіз проекту з точки зору цілей і вимог і на підставі цього формується приблизний план робіт з даними;

- *Попередній збір і аналіз даних.* У цій фазі виконуються попередні «розвідувальні» дії з аналізу даних. Фактично при цьому моделюються різні варіанти подальших дій по проекту і вибирається найбільш ефективний. Вибірка даних для цього етапу повинна бути невеликою, але достатньо представницькою;

- *Збір даних.* У цій фазі проекту на підставі обраної раніше стратегії його реалізації відбувається формування і накопичення необхідної бази початкових даних у необробленому вигляді;

- *Підготовка даних.* Включає дії з попередньої обробки даних і перетворення їх до виду, найбільш зручному для подальшої обробки. Необхідність даного етапу обумовлено тим, що початкові дані, що надходять з різних джерел, можуть бути зашумлені, неповні, мати різні формати і одиниці вимірювання. Крім того, ефективність деяких методів інтелектуального аналізу залежить від того, в якому вигляді представлені вхідні дані [61, 232];

- *Моделювання.* У цій фазі відбувається основна частина процесів аналізу

і обробки даних з метою отримання оптимальної моделі рішення задачі;

– *Оцінка результатів*. Включає ретельний розгляд отриманих результатів і понесених витрат, за підсумками якого приймається рішення про те, чи відповідають досягнуті результати цілям проекту та про доцільність їх впровадження;

– *Впровадження*. Включає інтеграцію розроблених моделей і методів в практичну діяльність. Залежно від особливостей задачі, що вирішується, етап впровадження може мати різну складність. Так, може знадобитися реалізація постійного автоматичного поновлення розроблених моделей на основі нових даних.

Слід зазначити, що послідовність фаз і їх набір не слід розглядати як жорстко заданий. При реалізації різних проектів він може змінюватися і доповнюватися. Так, якщо задачі, що вирішуються, носять постійний характер, то доцільно додати до вказаного переліку фазу *реалізації* проекту, тобто переведення комплексу розроблених моделей в систему програмного забезпечення, інтегровану з корпоративною інформаційною системою підприємства, або організації.

Наступний класифікаційний підхід має в своїй основі поділ задач в залежності від *типу даних*. Алгоритми аналізу і обробки різних типів даних можуть істотно відрізнятися, що обумовлює практичну значущість даного підходу. До основних типів даних в контексті цієї класифікації слід віднести [49]:

- числові;
- текстові;
- звукові;
- статичні зображення;
- динамічні зображення (відео);

Окрім того, в рамках кожного з зазначених класів може існувати розподіл на підкласи. Так, числові дані можуть бути представлені у вигляді рядів, матриць, або графів.

Нарешті, в залежності від мети виконуваних дій слід виділити задачі *аналізу* і *обробки даних*. Оскільки часто-густо ці терміни вживають як синоніми, все ж слід визначити відмінності між ними.

Визначення 1.1.

Аналізом даних будемо називати процес вилучення з необроблених даних відомостей, корисних для дослідника.

Це визначення ґрунтується на прийнятому в [35] підході до аналізу даних, як до процесу перекладу даних в інформацію. До такої інформації належать приховані в масиві даних зв'язки між різними його підмножинами. Знання про них дозволяє більш ефективно використовувати дані: робити більш достовірні прогнози, поліпшити якість пошуку інформації, підвищити ефективність діагностики і так далі.

Чіткого визначення терміну «*обробка даних*» у вітчизняній літературі немає. Аналіз [61, 232, 110, 46, 195 та ін.] дозволяє зробити висновок про необхідність тлумачення даного терміну в широкому і вузькому сенсах і дати

такі визначення.

Визначення 1.2.

Обробкою даних в широкому сенсі назвемо процеси, пов'язані зі збором, зберіганням, аналізом та трансформацією даних.

Визначення 1.3.

Обробкою даних у вузькому сенсі назвемо процеси, в результаті яких з одного масиву даних виходить інший, із заданими властивостями.

Визначенню 1.3 близько відповідає англomовний термін «*Data processing*», який можна трактувати як *перетворення даних*, метою якого є «вдосконалення» даних в рамках заданих критеріїв.

До обробки даних у вузькому сенсі можна віднести задачі оптимізації, фільтрації, ранжирування і інші, результатом яких є трансформація даних.

Дуальність тлумачення привела до того, що терміни «аналіз даних» і «обробка даних» часто використовуються як синоніми. Тому, щоб уникнути можливих різночитань, в даній роботі термін «обробка даних» буде використовуватися виключно у вузькому сенсі (за визначенням 1.3, якщо не вказано інше).

Слід зазначити, що при аналізі, розмірність даних на виході зазвичай не залежить від розмірності вхідних вибірки. У той же час при обробці даних розмірності вхідній та вихідній вибірок тісно пов'язані.

Залежно від умов конкретної економічної задачі процеси аналізу і обробки даних можуть слідувати в будь-якому порядку, а також здійснюватися незалежно один від одного. Окремо слід виділити випадок отримання метаданих, тобто частково структурованої інформації відносно неструктурованих даних.

Найбільш загальна класифікація задач *аналізу даних* виділяє в якості класифікаційної ознаки *мету вирішення*. Відповідно до неї ці задачі поділять на дві великі групи: *прогностичні* і *описові*.

Перші пов'язані з побудовою моделі, яка може використовуватися для прогнозування результатів поведінки аналізованої системи в тих випадках, про яких даних поки немає. До них відносяться задачі передбачення банкрутства, прогнозування фінансових часових рядів і багато інших.

Метою вирішення завдань другої групи є пошук і опис прихованих закономірностей в даних і виведення семантичних правил, які можуть використовуватися в подальшому для підвищення ефективності роботи. Тому дану групу також називають задачами структурного аналізу даних (*structured data mining*). До них відноситься велика кількість маркетингових задач, пов'язаних з аналізом різних цільових груп і виявленням схильностей їх учасників, задачі аналізу поведінки та інші.

Однак розглянуті групи задач не слід розглядати ізольовано одна від одної. Так, після дослідження закономірностей поведінки деякої системи в рамках рішення описової задачі, результати можна використовувати зокрема для прогнозування її поведінки.

В рамках методології моделювання інноваційних інтелектуальних систем прийняття рішень в економіці ця класифікація повинна бути розглянута більш

докладно. Узагальнення досліджень [49, 64, 62, 35, 17] дозволило скласти наступну класифікацію задач з аналізу даних (рис. 1.5).

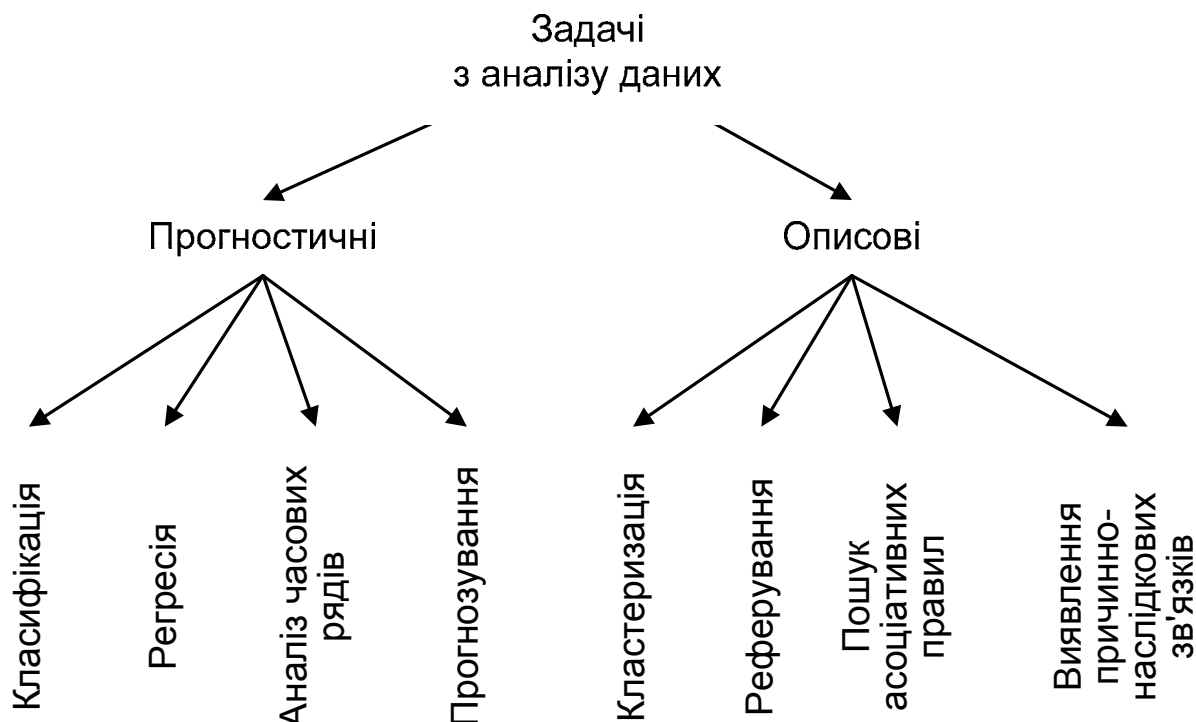


Рисунок 1.5 – Загальноприйнята класифікація задач аналізу даних

Розглянемо задачі, наведені на рис. 1.5.

Класифікація. Є одним з найбільш поширених прикладів прогностичних задач інтелектуального аналізу даних. В рамках її вирішується задача розділення деякої множини даних на заздалегідь визначений набір класів. Класи визначаються до обробки множини. Рішення задачі класифікації дозволяє користувачеві не тільки вивчити і проаналізувати існуючі дані, але і дає можливість передбачити майбутню поведінку даної вибірки [137].

За кількістю класів, на які поділяється вхідні вибірка, слід виділяти задачі *бінарної* та *полінарної* класифікації. При бінарній класифікації вхідна вибірка ділиться тільки на два класи (точніше - 1 клас і все інше). Незважаючи на таке просте трактування, до бінарної класифікації зводиться велика кількість економічних задач прийняття рішень, пов'язаних із вибором з двох альтернатив, одна з яких зазвичай трактується, як позитивна, а інша - як негативна. Наприклад - видавати кредит, або ні; прийняти, або відхилити проект; надійний, чи не надійний контрагент. До завдань бінарної класифікації можна звести також виявлення випадків шахрайства по кредитних картах, фільтрація електронної пошти та інші.

При полінарній класифікації вхідна вибірка ділиться на три, чи більше класи. Прикладами таких задач є: розпізнавання образів, визначення класу позичальника, сегментація клієнтів (якщо класи задані заздалегідь).

Незважаючи на схожість постановок задач бінарної та полінарної класифікації, методи і інструменти їх вирішення суттєво різняться.

Регресія. Це один з поширених класів задач прогностичної групи. Рішення регресійної задачі передбачає виявлення взаємозв'язку між незалежними (вхідними) змінними і залежними (вихідними). В результаті цього евристичним або аналітичним шляхом визначається деяка математична формула, що виражає залежність між вхідними та вихідними даними [137].

Прикладом задачі регресії може бути визначення суми кредиту, який може бути виданий клієнту.

Прогнозування є ще одним різновидом задачі регресії. Метою її рішення є наближене визначення значень деяких показників в майбутньому на підставі відомих значень в минулому і сьогоденні.

Прикладами задач цього класу є прогнозування зростання організації, прогнозування виконання бюджету, прогнозування потреб у залученні нових співробітників і тому подібні.

Аналіз часових рядів. Метою рішення цієї задачі є прогнозування майбутніх значень деякого набору даних, де значення вихідної змінної залежить не тільки від її минулих значень, але і від часу. Характерною ознакою часових рядів є рівномірний розподіл вхідних даних по годинах, днях, тижнях, і тому подібних одиницях виміру часу. Аналіз часових рядів є різновидом задачі регресії, але оскільки використовує специфічні вхідні дані і методи вирішення, виділяється в окремий клас.

До завдань даного класу можна віднести, наприклад, вивчення тенденцій на фондовому ринку для прогнозування курсових коливань.

Кластеризація відноситься до описової групі задач інтелектуального аналізу даних. При її вирішенні потрібно знайти закономірності в масиві даних, виділити в ньому деяку кількість зон (кластерів) і розподілити по ним дані. Задача кластеризації нагадує задачу класифікації, з тією суттєвою різницею, що самі класи заздалегідь не визначені. Для вирішення задач кластеризації використовуються алгоритми навчання без вчителя.

Прикладом кластеризації є угруповання потенційних споживачів товару в контекстній рекламі та інших різновидах маркетингу. Іншим прикладом є угруповання банків в банківській системі країни для подальшого аналізу їх стійкості [149].

Реферування. Задача автоматичного реферування набула особливої актуальності у зв'язку з розвитком Internet і різким зростанням обсягів інформаційних ресурсів. Суть її зводиться до створення на підставі вхідного документа (частіше за все – текстового) деякої вибірки, яка б при заданих обмеженнях на обсяг найбільш повно представляла його суть. Прикладом може служити автоматичне складання анотацій, або коротких описів статей для видачі їх в результатах роботи пошукових систем.

Спочатку задача автоматичного реферування вирішувалася виключно для тестових документів, проте в даний час сфера її застосування розширилася і охоплює всі основні види представлення інформації (зображення, звук, відео). Оскільки термін «реферування» в українській мові означає виключно обробку текстової інформації [73] (на відміну від вихідного англійського терміна *summarization*, сенс якого більш широкий), в даний час цю задачу доцільно

визначити, як «виділення значущих ознак». Серед актуальних прикладів її застосування – тематичний пошук контенту; контекстний пошук, виявлення матеріалів, що порушують інтереси правовласників, чи законодавчі обмеження і подібні.

Пошук асоціативних правил. Пошук асоціативних правил дозволяє встановити зв'язки і відносини між змінними у великих базах даних. Асоціативні правила дозволяють знаходити закономірності між пов'язаними подіями, тобто дають можливість відповісти на питання: «З якою ймовірністю пов'язані події А і Б?». При цьому послідовність настання подій значення не має.

Класичним прикладом застосування алгоритмів пошуку асоціативних правил є аналіз кошика покупок. Так, часто купуються разом хліб і молоко; текіла і лимон; шампанське, цукерки і презервативи. Це дає можливість раціонально розташувати товари в магазині, що збільшує його пропускну здатність і оборот. Методи пошуку асоціативних правил знаходять застосування зокрема в економічних задачах всіх різновидів.

Алгоритми пошуку асоціативних правил відносяться до класу навчання з вчителем.

Виявлення причинно-наслідкових зв'язків. Предметом задачі в цьому випадку є пошук статистично значущих закономірностей в послідовності даних. Тому в англійській літературі цей клас задач носить назву *Sequential pattern mining* (див., наприклад [24]).

На відміну від пошуку асоціативних правил, виявлення причинно-наслідкових зв'язків враховує фактор часу і дозволяє відповісти на питання: «З якою ймовірністю настання події А тягне за собою настання події Б?». Зазвичай передбачається, що події описуються дискретними величинами, що відрізняє дану задачу від аналізу часових рядів.

Виявлення причинно-наслідкових зв'язків є актуальним в системах інформаційної та фінансової безпеки. Зокрема, автоматичні системи безпеки банківських платіжних карт орієнтовані на виявлення специфічних дій, характерних для шахрайських операцій. Також дана задача вирішується в медицині при виявленні причин виникнення захворювань, а також у багатьох інших.

Класифікація, що наведена на рис. 1.5, охоплює широкий спектр актуальних економічних завдань. Однак ускладнення соціально-економічних процесів і методів їх аналізу постійно викликають появу нових задач. Так у фундаментальному дослідженні *Jiawei Han* та *Micheline Kamber* по концепціям і технологіям аналізу даних розглядаються такі задачі, як виявлення фрагментів в послідовності даних (пошук підпослідовностей) [24, с 248], пошук аномалій [24, с. 543], аналіз соціальних мереж [23, с 556].

Задача пошуку підпослідовностей може розглядатися, як різновид задачі кластеризації, стосовно аналізу динамічних, або часових рядів. В рамках її вирішення потрібно виділити послідовність даних (фрагмент), відповідний іншій послідовності (запиту). При цьому розмірність запиту і фрагмента (кількість елементів ряду, що входять в них) може мати відчутні відмінності.

Різновидом цієї ж задачі є пошук *фракталів*, тобто послідовностей даних, які мають властивість самоподібності (точного або приблизного збігу даних зі своєю частиною). Фрактали можуть виникати в послідовності, що описують колективну поведінку людей (ефект натовпу).

Дана задача виникає, наприклад, при аналізі біржових даних, аналізі продажів.

Пошук аномалій і викидів. Включає широкий спектр задач, пов'язаних з виявленням та ідентифікацією таких елементів даних, які не відповідають виявленим раніше закономірностям. Причому такі аномалії можуть бути як новими закономірностями (наприклад, нові тренди в біржових даних), так і сигналами про ненормальну поведінку об'єкта спостереження.

Об'єктом пошуку аномалій можуть бути результати вимірювань, часові ряди, текстова інформація, графи.

Задача пошуку аномалій вирішуються в рамках систем інформаційної безпеки, систем виявлення шахрайських операцій з банківськими картами, медичної діагностики, перевірки правопису і грамотності і в багатьох інших областях. Крім того алгоритми пошуку аномалій може використовуватися при попередній обробці даних для їх очищення від шумів.

Методи, які використовуються для пошуку аномалій, можна розділити на три групи [24]:

1. Методи контрольованих алгоритмів (навчання з учителем). В цьому випадку необхідно мати набір даних розмічених на «нормальні» і «аномальні»;
2. Методи неконтрольованих алгоритмів (навчання без вчителя). Припускають наявність нерозміченого набору даних, в якому алгоритм виділяє закономірності і відхилення від них;
3. Методи частково-контрольованих алгоритмів. Передбачається наявність при навчанні тільки даних, розмічених як «нормальні».

За своєю суттю задача пошуку аномалій і викидів близька до задачі виділення значимих ознак, що дає підстави об'єднати їх в одну класифікаційну групу (таксон) – *аналіз відхилень*.

Аналіз соціальних мереж - це процес дослідження соціальних структур, за допомогою представлення їх у вигляді вершин (вузлів) і зв'язків між ними (ребер), з використанням інструментарію графів і теорії мереж.

Класичні методи аналізу дозволяють досить ефективно досліджувати топологію соціальних мереж, зокрема визначати ключові вузли і ребра, виділяти слабкі ланки, знаходити найкоротші шляхи між вузлами. Однак в сучасних умовах цього недостатньо, оскільки роль соціальних мереж безупинно зростає, а, отже, зростає і складність запитів до них.

Так, в ході президентської кампанії 2016 року в США перемога була здобута кандидатом, який більш ефективно використовував соціальні мережі. З них витягалася інформація про потенційних виборців, їхніх інтересах, зв'язках, реакціях на події. Все це використовувалося для проведення найбільш ефективних «точкових» впливів, з індивідуальним підходом до кожного виборця [1].

Значна кількість практичних завдань з аналізу соціальних мереж

зводиться до виділення в мережах груп учасників за деякою узагальнюючою ознакою. Чим більш абстрактним є ця ознака, тим складнішим стає вирішення задачі.

Задачі аналізу соціальних мереж є актуальними в економіці, маркетингу, соціології, політиці, історичних дослідженнях і багатьох інших областях.

Вивчення і зіставлення розглянутих задач з аналізу даних дозволило актуалізувати та узагальнити класифікацію, наведену на рис. 1.5. Уточнена класифікація показана на рис. 1.6.

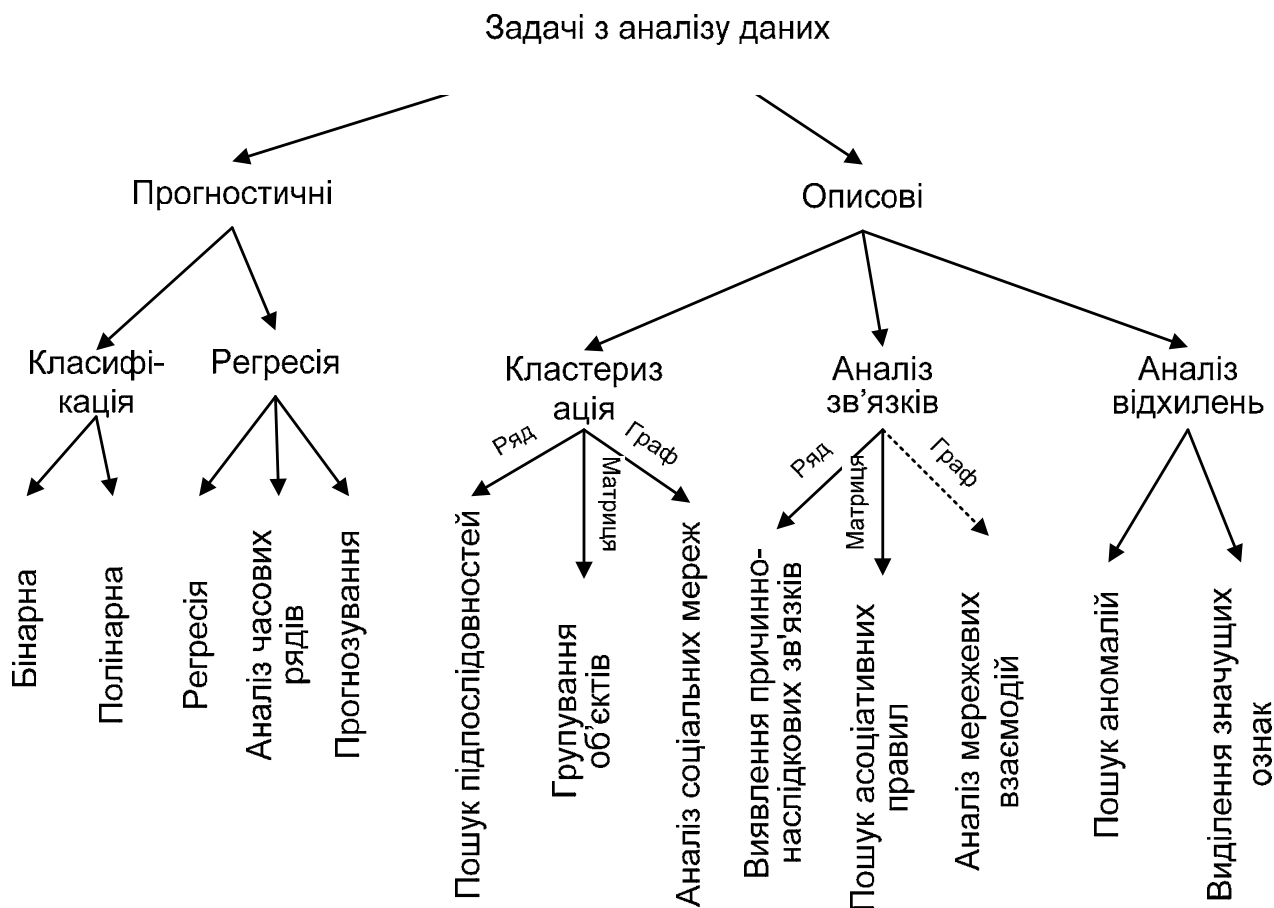


Рисунок 1.6 – Уточнена класифікація задач аналізу даних

Як видно з порівняння рис. 1.5 і рис. 1.6, кількість рангів класифікації збільшено, за рахунок виділення і впорядкування таксономічних ознак. Це дозволяє краще зрозуміти взаємозв'язок між різними класами задач і методами їх вирішення, що підвищує ефективність аналізу даних.

Розглянемо особливості запропонованої класифікації.

До таксону *регресії* віднесено задачі, що потребують виявлення взаємозв'язку між вхідними та вихідними змінними. За таким визначенням, задачі прогнозування та аналізу часових рядів також слід віднести до регресійних. Відмінності між ними полягають у розташуванні вхідних та вихідних змінних у часі. Метою *прогнозування* є наближена оцінка значень деяких показників в майбутньому на підставі відомих значень в минулому і

сьогоденні. Метою *аналізу часових рядів* є прогнозування майбутніх значень деякого набору даних, де значення вихідної змінної залежить не тільки від її минулих значень, але і від часу.

Розглянемо таксономічну категорію задач кластеризації. Проведений аналіз показав, що задачі пошуку підпоследовностей в даних і аналізу соціальних мереж фактично також є задачами кластеризації, відрізняючись тільки представленням вхідних даних і розмірністю простору угруповання.

Так, пошук підпоследовностей в динамічних рядах фактично являє собою *задачу кластеризації в одновимірному просторі* часу, де кожен елемент має тільки двох сусідів – попереднє значення ряду і наступне.

Кластеризацією в чистому вигляді зазвичай називають задачі угруповання масиву вхідних даних, де кожен елемент (крім крайніх і кутових) має *фіксоване число сусідів*. У двовірному просторі їх чотири, в тривірному – шість, і так далі. Вхідні дані в такому випадку представляються в матричному вигляді.

Нарешті, аналіз соціальних мереж – це задача, де кожен елемент вхідних даних (вершина графа) може мати *будь-яку кількість сусідів*, яке в реальних соціальних мережах може досягати декількох тисяч (а для окремих людей і більше).

З цього можна зробити висновок, що саме ознака розмірності простору угруповання (кількості сусідніх елементів у вхідних даних) є природною класифікаційною ознакою для задач кластеризації, що і відображено на рис. 1.6.

Задачі *пошуку асоціативних правил і виявлення причинно-наслідкових зв'язків* розміщено під таксоном «Аналіз зв'язків». Підставою для цього є схожість завдань. Як і для таксона «Кластеризація» основна відмінність між ними полягає в способах представлення вхідних даних і розмірності їх простору.

При виявленні причинно-наслідкових взаємозв'язків, події (аналізовані елементи) поділяються за часом і тому можуть бути представлені у вигляді динамічного ряду.

При пошуку асоціативних правил вхідний набір даних представляється в матричному вигляді, що описує події, які відбуваються одночасно, або близькі за часом.

Оскільки структура отриманого таксона «Аналіз зв'язків» нагадує структуру таксона «Кластеризація», це дає підстави припустити існування ще однієї задачі, що раніше не розглядалася, а саме задачі аналізу зв'язків, вхідні дані для якої представлені у вигляді графа.

Якщо сформулювати її суть, за аналогією з іншими задачами цього таксона, то можна припустити, що повинен здійснюватися пошук епіцентрів мережевої активності, тобто вузлів, що впливають на процеси, які відбуваються в мережі, або ініціюють такі процеси. У широкому сенсі ця задача зводиться до виявлення для кожного вузла мережі *керуючих і керованих вузлів*. Таким чином, її можна сформулювати, як *аналіз мережевих взаємодій*.

В даний час схожа задача – пошук «нульового пацієнта», тобто людини, з якої розпочинається епідемія тяжких захворювань, існує в медицині. Проте

методи її вирішення передбачають тільки ретроспективний аналіз і не можуть використовуватися для передбачення того, хто може стати таким «нульовим пацієнтом» в майбутньому.

В економіці аналіз взаємодій може використовуватися для прогнозування поширення кризових явищ, підвищення ефективності рекламних кампаній, виявлення цільових клієнтів.

Таким чином, уточнена класифікація задач аналізу даних (рис. 1.6) дає можливість не тільки показати взаємозв'язок між різними класами задач і методами їх вирішення, а й виявити критерій розподілу задач на групи в рамках окремих таксонів (спосіб представлення вхідних даних у вигляді ряду, матриці, або графа), а також сформулювати завдання, які раніше не виділялися в окремий клас.

Розглянемо далі класифікацію задач *обробки даних*.

Обробці як правило підлягають дані, що є первинною інформацією, отриманою за місцем виникнення. Обробка даних може бути як підготовчим етапом для подальшого аналізу, так і самостійним процесом.

І у вітчизняній і в зарубіжній літературі відсутня чітка класифікація задач обробки даних. Але, з позицій запропонованого вище визначення 1.3 до таких задач можна віднести сортування, фільтрацію, ранжирування, відновлення, очищення і квантування даних.

Сортування - задача, пов'язана з розташуванням елементів даних в заданій послідовності, або розподілом їх по групах. При наявності єдиного критерію сортування, або критеріїв, які перебувають в строгому ієрархічному підпорядкуванні, задача сортування даних має тривіальне рішення. Складність її різко зростає з ускладненням структури поля критеріїв і його визначеності. Так, задача сортування товарів в порядку зростання, або зменшення клієнтських переваг не має тривіального рішення.

Термін *фільтрація* має велику кількість значень, стосовно до різних областей знань, в яких використовується (наприклад, у фізиці, хімії, математики, радіотехніці, обробці зображень). При обробці економічних даних під фільтрацією будемо мати на увазі вибірку інформації, що задовольняє заданому критерію. Вхідні дані для фільтрації можуть бути представлені у вигляді ряду, масиву, або графа. Складність фільтрації зростає з ускладненням структури поля критеріїв і зменшенням їх визначеності.

Задача *ранжирування* відрізняється від сортування тим, що для її вирішення необхідно визначити метод визначення рангу кожного елемента вхідної послідовності даних, тобто метод, що дозволяє згорнути весь вектор значень елемента в єдиний параметр - ранг. Це дозволяє порівнювати між собою будь-яку кількість елементів з вхідного набору, не вдаючись до його повного сортування, що ефективно при наявності великих обсягів даних для обробки. Складність ранжирування зростає з ускладненням структури поля критеріїв. Наприклад, ранжирування прострочених кредитів за ступенем ймовірності їх повернення не може мати однозначного рішення, оскільки методи визначення цієї ймовірності самі по собі засновані на припущеннях.

Необхідність у *відновленні* даних виникає в тому випадку, якщо вхідна

вибірка містить пропуски, або якісь дані в ній відсутні, але є гіпотези про природу їх виникнення, що дозволяє оцінити найімовірніші значення. Відновлення даних має велике значення при використанні методів машинного навчання, оскільки дозволяє підвищити його ефективність за рахунок розширення вхідній вибірки даних. Необхідність відновлення даних може виникати при їхньому поданні в будь-який з розглянутих форм - у вигляді рядів, матриць, або графів. В останньому випадку об'єктом відновлення виступають відсутні у вхідній вибірці зв'язки між вершинами графа

Задача *очищення* даних, по суті, є зворотною, відносно задачі фільтрації і передбачає виключення з початкової вибірки «зайвих» даних. Очищення даних можна розглядати по відношенню до рядів, матриць і графів.

Очищення *рядів* передбачає усунення «викидів», тобто даних, які явно виходять за межі основної тенденції. Причинами появи таких викидів можуть бути як помилки вимірювання, так і навмисні спотворення інформації. У будь-якому випадку, спотворені дані роблять сильний вплив на якість подальшого аналізу, особливо при використанні виключно формальних методів, в тому числі машинного навчання.

Очищення даних, представлених в *матричній формі*, проводиться для усунення з вхідній вибірки даних факторів, що мало впливають на вихідні показники, або зовсім не пов'язані з ними. Ця процедура обов'язково повинна попереджувати аналіз даних.

Відносно до *графів* очищення передбачає проріджування зв'язків і застосовується як по відношенню до вхідних даних, так і по відношенню до деяких видів економіко-математичних моделей, структурою яких є граф, наприклад штучних нейронних мереж. При цьому усуваються слабкі зв'язки, мало впливають на загальний результат, але істотно ускладнюють аналіз.

Квантування використовується по відношенню до динамічних рядів даних при необхідності зменшення кількості їх елементів, тобто при спрощенні рядів. Існують такі види, як квантування за рівнем і квантування за часом.

У першому випадку діапазон значень неперервної або дискретної величини розбивається на кінцеве число інтервалів. Якщо протягом деякого періоду часу значення величин динамічного ряду не виходило за межі одного інтервалу, то в результаті квантування всі ці величини будуть замінені одним значенням [109, с. 184].

У другому випадку розбиття динамічного ряду на інтервали відбувається за часом. Значення, які брали елементи ряду в межах кожного інтервалу, замінюються одним усередненим, або декількома, що відображають граничні значення показника за період. Останній спосіб використовується для представлення біржових даних.

Існують також різновиди квантування за рівнем і за часом зі змінним кроком квантування [155].

Зазначимо, що задача квантування формулюється в даний час виключно по відношенню до динамічних рядів даних. Однак її знаходження в одному таксоні з задачами, які можуть вирішуватися по відношенню до матриць і графів дає підставу припустити можливість розглядати задачу квантування і по

відношенню до інших форм представлення даних. Зокрема, такі інструменти обробки даних, як метод головних компонент, а також згорткові штучні нейронні мережі фактично виконують квантування даних, які представлені у вигляді матриць [243].

Таким чином, категорія *обробка даних* об'єднує широкий спектр задач, які за складністю варіюються від тривіальних до таких, для яких отримання точного рішення практично неможливо. З огляду на те, що останні можуть зустрічатися в кожному з розглянутих класів, можна говорити про задачі інтелектуальної обробки даних, класифікацію яких представимо в такий спосіб (рис. 1.7).

Як видно з аналізу рис. 1.7, основною класифікаційною ознакою задач інтелектуальної обробки даних пропонується ознака впливу на порядок елементів вхідної вибірки даних.

Зміна порядку елементів відбувається при вирішенні задач ранжирування і сортування. Вхідні дані при цьому зазвичай представлені у вигляді рядів, або зводяться до них. В інших розглянутих задачах зміна порядку елементів вхідних даних не відбувається.

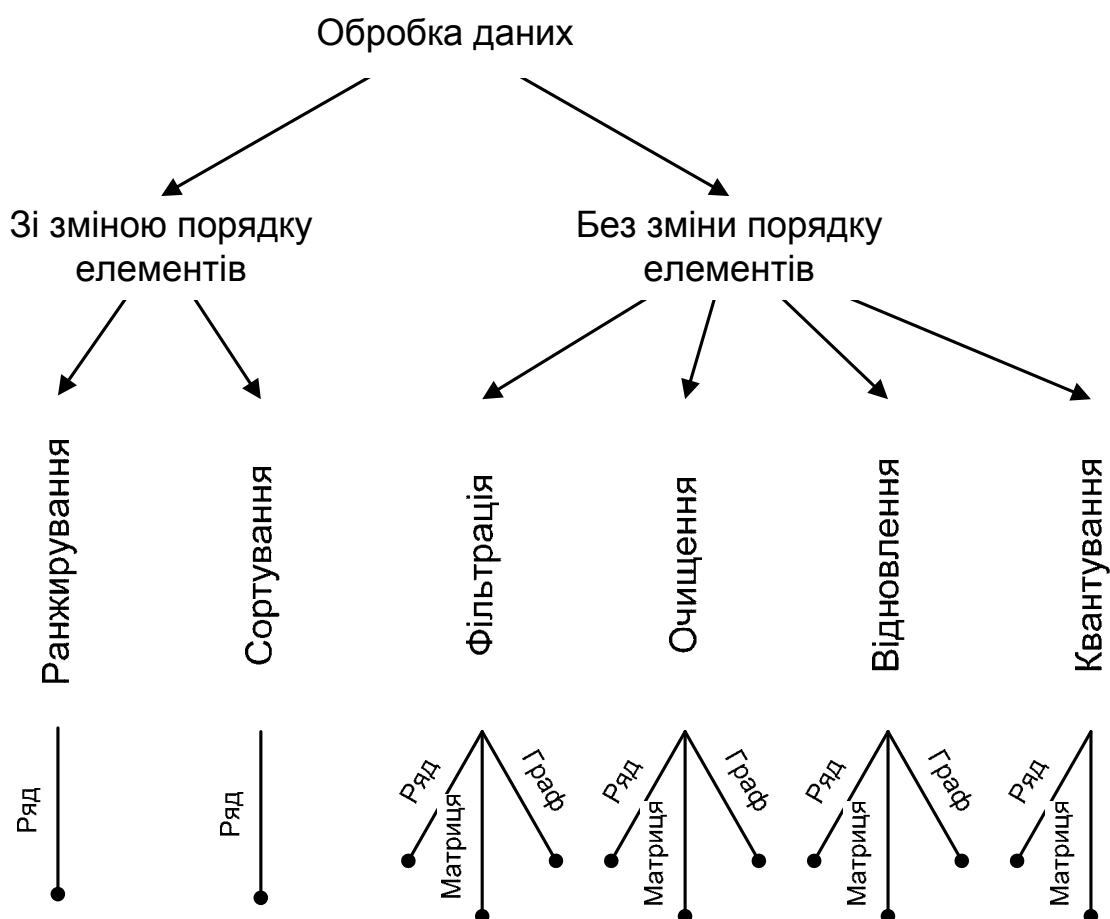


Рисунок 1.7 – Класифікація задач інтелектуальної обробки даних

При цьому слід зазначити, що методи розв'язання задач по відношенню до даних, які представлені у вигляді рядів, матриць і графів істотно розрізняються.

Таким чином, проведений аналіз задач аналізу і обробки даних дозволив провести їх класифікацію та розподіл за основними таксонами. Слід зазначити, що, з огляду на динамічний розвиток інтелектуальних методів прийняття рішень, формування повної і несуперечливої класифікації задач аналізу і обробки даних не представляється можливим. Тому отримані результати слід розглядати як досить правдоподібні гіпотези про природу цих задач.

Розглянемо далі класифікацію методів, які використовуються для вирішення задач інтелектуального аналізу і обробки даних – інтелектуальних обчислень.

Визначення 1.4.

Інтелектуальними обчисленнями будемо називати методи та системи штучного інтелекту, які спрямовані на підтримку прийняття рішень, зокрема на вирішення задач інтелектуального аналізу даних, обробки даних, оптимізації.

Можна виділити два класичних підходу до класифікації та розробки методів інтелектуальних обчислень – «нисхідний» і «висхідний» [12].

Нисхідний, або семіотичний підхід (англ. *Top-Down AI*) має на увазі створення систем штучного інтелекту на основі імітації високорівневих психічних процесів, таких як мислення, міркування. Результатом застосування цього підходу є, наприклад, експертні системи, бази знань, системи логічного висновку (включаючи нечітку логіку).

Висхідний, або біологічний підхід (англ. *Bottom-Up AI*) має на увазі створення систем штучного інтелекту на основі моделювання базових біологічних і фізичних процесів. Результатами застосування цього підходу є такі інструменти штучного інтелекту, як штучні нейронні мережі.

Подальший розвиток методів аналізу даних змусив розширити цю класифікацію. Різні школи пропонують різні її варіанти, але в рамках даного дослідження будемо виділяти в якості самостійних агентно-еволюційний та імітаційний підходи до проектування інтелектуальних систем.

Визначення 1.5.

Агентно-еволюційний підхід має на увазі використання для рішення задачі деякого набору самостійних програм - агентів. Агенти діють в програмно-створеному середовищі, характеристики якого залежать від умов розв'язуваної задачі. Кожен агент наділений деякими можливостями щодо сприйняття умов середовища, і вибору варіантів дій, виходячи з цих умов. Для даного підходу характерна не тільки часова, а й просторова неоднорідність створюваної моделі. В рамках агентно-еволюційного підходу можна виділити генетичні алгоритми та інші непереборні методи вирішення NP-повних задач (імітація відпалювання, метод мурашиних колоній)

Визначення 1.6.

Імітаційний підхід дозволяє створювати і досліджувати моделі систем, для яких відомі лише деякі закономірності поведінки (так звана «чорна скриня»). Переваги імітаційного підходу виявляються тоді, коли необхідно отримати знання про поведінку системи, яка складається з безлічі «чорних скринь», що звичайними методами зробити неможливо.

Розглянемо основні особливості інтелектуальних засобів підтримки

прийняття рішень, які використовуються в даний час для рішення економічних задач.

Експертні системи.

Хоча теорія експертних систем в сучасному вигляді була розроблена в 1970-х роках, проте, сама концепція логічного вибору певного варіанту на основі аналізу деякого набору вихідних передумов набагато старше [127]. Практичне використання експертних систем почалося з появи скорингових моделей і методів оцінки потенційних кредитних ризиків [13].

Сучасна теорія експертних систем передбачає обов'язкову наявність бази знань, яка містить сукупність фактів і правил логічного висновку з них, якими користується система. Включення правил логічного висновку до складу баз знань докорінно відрізняє їх від баз даних, де міститься тільки інформація. Зокрема, система, логічного висновку може на підставі аналізу інформації отримувати знання, які в чистому вигляді були відсутні у вихідних базах.

Згідно Дж. Джарратано і Г. Райлі експертні системи можуть використовуватися для вирішення таких завдань, як інтерпретація даних; діагностика; моніторинг; проектування; прогнозування; зведене планування; оптимізація; навчання; керування; ремонт; налагодження [107].

Найважливішою з переваг експертних систем, як інструменту підтримки прийняття рішень, є використання концепції «білої скрині», тобто прозорість одержуваних висновків. Будь-який результат, отриманий за допомогою експертної системи, може бути пояснений на основі вхідних даних і наявних в системі знань про об'єкт.

Нечітка логіка.

Нечітка логіка, як інструмент підтримки прийняття рішень, використовується в умовах неповноти, нечіткості, або розмитості інформації про даний об'єкт. Нечіткий опис предметної області набагато ближче до природної мови і образу думок людини, ніж опис в термінах формальної логіки. Так, засобами нечіткої логіки можуть бути записані і використані для виведення такі відносини, як «можливо», «скоріше так, ніж ні» «іноді» і тому подібні. Це дозволяє ефективно використовувати даний інструментарій для опису предметної області та наповнення баз знань.

Математичний апарат нечіткої логіки дозволяє не тільки описувати такі елементи, але і чинити з ними операції, еквівалентні таким в класичній логіці, а також інтерпретувати результати з практичної точки зору. Основою нечіткої логіки є положення про те, що ступінь приналежності елемента до деякої множини може приймати значення не тільки «0», або «1», а й інші, в інтервалі $[0..1]$. При цьому будь-який об'єкт може одночасно належати кільком множинам. Так, людина у віці 35 років іншими людьми може сприйматися і як «молодий» і як «середнього віку», що не піддається опису методами класичної логіки, але легко описується в термінах логіки нечіткої.

Розробка будь-якої системи з використанням апарату нечіткої логіки включає мінімум три обов'язкові процедури [114]:

1. Фазифікація - визначення нечітких множин і правил перекладу вхідних даних у нечіткі (включаючи визначення лінгвістичних змінних і функцій

приналежності).

2. Рішення задачі в нечітких термінах – створення правил обробки отриманих даних.

3. Дефазифікація - визначення правил перекладу результатів рішення в чіткі величини, які можуть використовуватися в подальших процедурах.

З напрямків використання нечіткої логіки найбільшого поширення набули системи нечіткого логічного висновку, серед яких слід виділити алгоритм Мамдані [39], алгоритм Сугено [66, 65], алгоритм Цукамото [67], алгоритм Ларсена [36].

Слід зазначити, що нечітку логіку можна розглядати не тільки як самостійний інструмент вирішення економічних завдань. Будучи використана в якості логічного базису при побудові експертних систем, нейронних мереж та інших подібних систем, нечітка логіка виводить їх на якісно новий рівень, часто істотно покращуючи ефективність роботи. Цим пояснюється стійке зростання значення нечіткої логіки в системах підтримки прийняття рішень.

Дерева прийняття рішень.

Являє собою багаторівневу бінарну ієрархічну структуру, де кожен вузол (гілка) відповідає за вибір одного з двох варіантів. Підставою для вибору є значення з набору вхідних даних. Прохід по гілках продовжується до тих пір, поки не буде досягнутий один з варіантів вирішення – лист (рис. 1.8). При цьому дерево може тільки розгалужуватись. Сходження гілок не допускається.

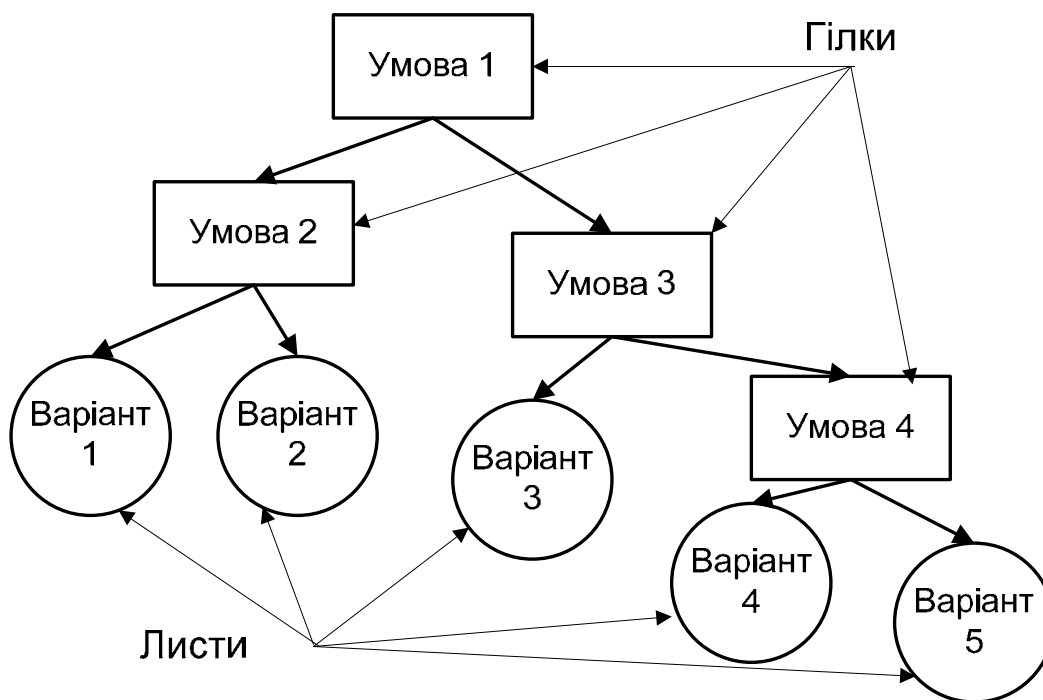


Рисунок 1.8 – Структура дерева прийняття рішень

Як видно з рис. 1.8, сутність роботи дерева прийняття рішень досить проста і полягає в багаторазовій перевірці умов «якщо-то». Тому відлік використання цього інструменту для вирішення практичних завдань слід вести не з його винаходу (встановити дату якого яку не представляється можливим), а з алгоритмів автоматичної побудови дерева. В даний час розроблено досить

велику кількість таких алгоритмів, серед яких:

CLS (Concept Learning System) [30]. Цей алгоритм є найбільш простим і ґрунтується на мінімізації суми вартості вимірювань (величини дерева) і вартості помилки класифікації.

ID3 (Iterative Dichotomiser 3) [56]. Даний алгоритм в якості основного критерію оптимальності використовує мінімізацію інформаційної ентропії.

Алгоритм ID3 має високу швидкодію і застосовується до цих пір, проте дерева, одержувані з його допомогою, мають зайву гіллястість.

C4.5 і C5.0 [55]. У них додано параметри управління гіллястістю дерева, процедури відсікання гілок, можливість роботи з числовими атрибутами, а також можливість побудови дерева з неповною навчальною вибіркою, в якій відсутні значення деяких атрибутів.

CART [5]. На відміну від перелічених вище алгоритмів, здатних вирішувати тільки задачі класифікації, цей може вирішувати також задачі регресії. Однак ефективність його роботи декілька нижче.

QUEST [37]. Один з нових алгоритмів, що дозволяє реалізувати багатовимірне розгалуження по лінійним комбінаціям порядкових предикторів і містить ряд засобів підвищення надійності і ефективності одержуваних дерев класифікації.

CTree (*Conditional inference trees*) [28]. Найсучасніший алгоритм, який використовує принцип рекурсивного розбиття на основі перестановок.

Порівняльні тести шести основних алгоритмів побудови дерев прийняття рішень, проведені в [63] показали лише незначну перевагу нових алгоритмів над C4.5, C5.0.

Дерева прийняття рішень можуть використовуватися для вирішення завдань опису даних; класифікації; регресії. Основним їх перевагою є наочність результатів. При цьому практично не потрібна підготовка даних. Метод може працювати одночасно з категоріальними і інтервальними змінними. Отримане дерево може бути використано для аналізу взаємозв'язків в даних і вилучення інформації про предметну область, тобто реалізується концепція «білої скрині».

В [156] показано, що дерева рішень можуть бути використані для аналізу значимості даних і відбору вхідних параметрів для нейронних мереж.

Штучні нейронні мережі

Поняття «штучної нейронної мережі» було запропоновано в 1943 році [44], а перша успішна модель - «Персептрон» реалізована в 1957 р [59] Концепція персептронних мереж до теперішнього часу залишається найбільш поширеною при вирішенні практичних завдань, включаючи економіку. Розглянемо її.

Персептронна ШНМ є двомірною структурою, яка складається з однотипних елементів – нейронів, згрупованих в шари (рис. 1.9).

У математичній моделі нейрона кожен синаптичний зв'язок характеризується деяким ваговим коефіцієнтом, визначеним в ході процесу навчання. Це дозволяє розглядати безліч синаптичних зв'язків як «пам'ять» ШНМ. Нейрон здійснює підсумовування вхідних сигналів і обробку їх активаційною функцією.

Приховані шари здійснюють основну обробку даних, забезпечуючи нелінійність реалізованих функцій. Незважаючи на те, що доведено принципову достатність одного прихованого шару для апроксимації функції будь-якої складності [232], іноді кілька прихованих шарів дозволяють отримати вигоду в кількості нейронів і швидкості навчання мережі.

Можливість навчання є головною особливістю нейронних мереж. Навчання ШНМ проводиться за допомогою спеціалізованих алгоритмів, що виконують настройку вагових параметрів синаптичних зв'язків з масиву вхідних даних. В результаті ШНМ при правильному підході до їх побудови і навчання отримують можливість вірно розпізнавати навіть ті об'єкти, яких не було в навчальній вибірці.

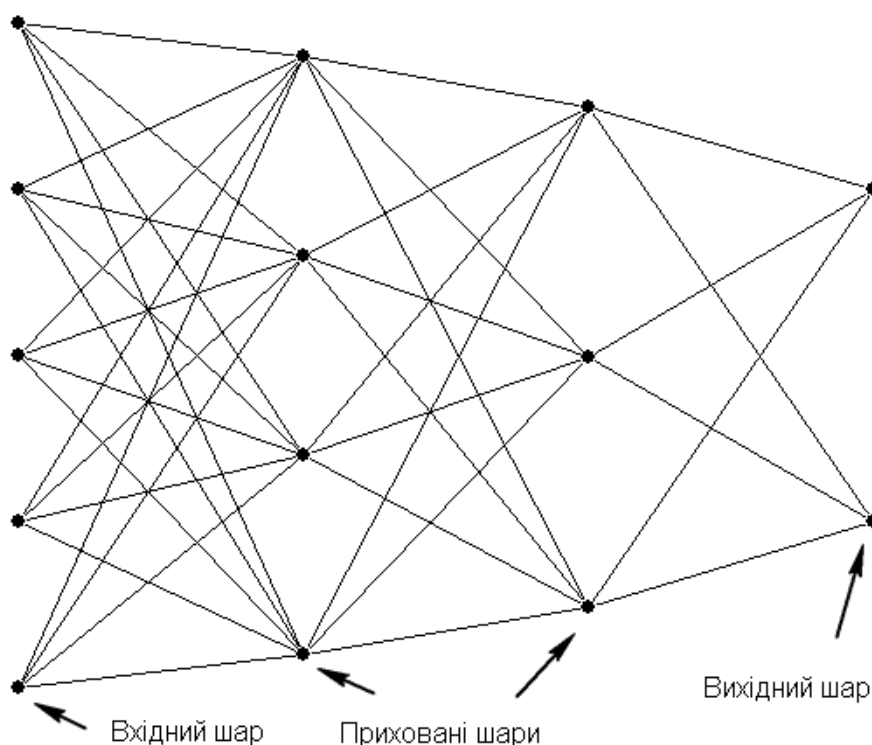


Рисунок 1.9 – Структура персептронної штучної нейронної мережі

В даний час для нейронних мереж можлива реалізація наступних базових парадигм машинного навчання [232]:

навчання з учителем – нейронна мережа налаштовується на запам'ятовування і відтворення залежностей між відомими станами входів і виходів. Безліч таких станів називається навчальною вибіркою. Чим більше кількість входів і складніше залежність, тим більше прикладів має бути в навчальній вибірці. Навчання з вчителем зазвичай має на увазі рішення задач регресії, або класифікації.

навчання без вчителя – в цьому випадку навчальна вибірка містить лише вхідні дані. Навчання має налаштувати нейронну мережу на пошук внутрішніх закономірностей в даних. При цьому зазвичай вирішується задача кластеризації вхідної множини. Прикладом ШНМ, орієнтованих на навчання без вчителя, є самоорганізаційні нейронні мережі, або мережі Кохонена [33];

навчання з підкріпленням – в процес навчання вводиться додаткова система – зовнішнє середовище, яке формує відгуки на рішення, прийняті нейронною мережею і таким чином виступає в ролі вчителя. Навчання з підкріпленням займає проміжне місце між попередніми парадигмами;

До 1986 р були відомі алгоритми навчання тільки для нейронних мереж з лінійної активаційною функцією, що істотно обмежувало застосування ШНМ. Розробка в зазначеному році алгоритму зворотного поширення помилки *Back Propagation* [60] дозволила вивести нейромережеві системи з кризи і сприяла сплеску інтересу до цього методу аналізу даних, а також масовому впровадженню ШНМ в практичних розробках.

Удосконалення алгоритмів навчання нейронних мереж триває до теперішнього часу. До сучасних крупних досягнень можна віднести створення в 2007 р алгоритмів глибинного навчання, розрахованих на роботу з великими багат шаровими нейронними мережами в програмах по реалізації повноцінного штучного інтелекту (розпізнавання образів, моделювання свідомості).

Нейронні мережі в даний час є найбільш поширеним і найбільш універсальним інструментом інтелектуального аналізу даних. Відомо більше 20 типів нейронних мереж, різноманітність яких така велика, що один з найбільших вчених - систематизатор в цій області, Саймон Хайкін, пропонує для них більш загальний термін - «самоорганізаційні системи, здатні до навчання» [232, с. 691]. В сукупності вони дозволяють вирішити практично будь-яку задачу з аналізу даних, зокрема розпізнавання образів і класифікація, прийняття рішень і управління, кластеризація, прогнозування, апроксимація, стиснення даних і асоціативна пам'ять, аналіз даних, оптимізація.

Таким чином, нейронні мережі мають такі переваги, порівняно з іншими інтелектуальними методами аналізу даних:

- *універсальність* – нейронні мережі можуть бути використані при вирішенні широкого спектра задач;
- *робота з інформацією будь-якої природи* – в тому числі з графічною інформацією, нечіткими даними;
- *ефективні методи контролю процесу навчання* – по динаміці процесу навчання можна робити висновки про його якість;
- *поширеність* – великий вибір програмних продуктів, що дозволяють створювати нейромережеві моделі;
- *толерантність до помилок* – навіть при недосконалому проектуванні нейромережевої системи, або значних шумах і помилках у вхідних даних, можна отримати позитивний результат;
- *можливість апаратної реалізації* – при необхідності та за умов достатнього фінансування це дозволяє отримати багаторазового збільшення швидкодії.

Разом з тим, ШНМ, як інструмент підтримки прийняття рішень, має і деякі недоліки:

- *концепція «чорної скрині»* – аналіз узагальнень, які отримала ШНМ в процесі навчання неможливий, або вкрай утруднений;
- *високі вимоги до продуктивності* обчислювальних систем – при

програмної реалізації великих ШНМ доводиться організовувати розподілені обчислення на великій кількості комп'ютерів;

– *високий рівень емпірики* в застосуванні ШНМ – для багатьох параметрів нейромережових моделей (кількість шарів, кількість нейронів, обсяг навчальної вибірки) відсутні достовірні способи визначення. Різні шляхи вирішення однієї задачі призводять до істотної різниці в результатах [161];

– *проблема навчання ШНМ*, пов'язана з тим, що неможливо заздалегідь визначити момент зупинки процесу навчання, що забезпечує найкращі апроксимуючі властивості отриманої нейронної мережі. Після деякої кількості ітерацій навчального алгоритму якість розпізнавання прикладів на тестовій вибірці починає погіршуватися. Це пов'язано з тим, що ШНМ починає «запам'ятовувати» навчальні приклади замість того, щоб шукати залежності між ними і називається *ефектом перенавчання*.

Генетичні алгоритми

Даний метод є представником агентно-еволюційного підходу до створення інтелектуальних систем аналізу даних. Як і нейронні мережі, генетичні алгоритми (ГА) засновані на математичній інтерпретації процесів, що відбуваються в живій природі. Базовими роботами в теорії генетичних алгоритмів є праці Дж. Холланда про адаптацію в природних і штучних системах. [27, 26].

На відміну від нейронних мереж, генетичні алгоритми орієнтовані на задачі пошуку рішень в багатовимірних просторах, зокрема – задачі оптимізації. Тому області застосування цих інструментів практично не перетинаються, навпаки, використання їх комбінацій в ряді випадків дозволяє поліпшити ефективність вирішення економічних завдань [103, 196].

Генетичні алгоритми виділяються серед методів оптимізації досить високими швидкістю і якістю роботи, а також універсальністю. Так, до функції, максимум, або мінімум якої відстежується за допомогою ГА, не пред'являється абсолютно ніяких вимог. Вона може бути переривчастою, недиференційованою, або складатися з шматків, які описуються різними рівняннями. У будь-якому випадку, алгоритм дозволить знайти для неї оптимальне (або близьке до нього) значення. Незважаючи на те, що генетичні алгоритми не завжди дозволяють знаходити абсолютно краще рішення, вони гарантують, що рішення буде досить хорошим. ГА істотно (іноді - на 4 порядки) підвищують ефективність вирішення задач пошуку екстремуму, на відміну від історично сформованих підходів до вирішення завдань оптимізації: переборних методів і методів локальної оптимізації (зокрема - градієнтних) [233]. Досвід використання ГА для вирішення різних задач [137, 166] підтверджує ефективність цього методу.

Генетичні алгоритми можуть бути використані для вирішення будь-яких задач, які можна звести до оптимізаційних. Серед них оптимізація функцій, різноманітні задачі на графах, підбір параметрів моделей (в т.ч., дерев рішень, ШНМ), задачі компонування, складання розкладів, визначення оптимальних ігрових стратегій [180].

Інші методи оптимізації

Крім генетичних алгоритмів для пошуку оптимальних рішень в багатовимірних просторах можуть використовуватися також такі агентно-еволюційні методи, як алгоритм імітації відпалювання, метод мурашиних колоній, метод бджолиних колоній, метод рою часток і деякі інші [177]. В основі цих методів лежать біологічні, або фізичні процеси, пристосовані до вирішення завдань оптимізації

Основними недоліками, що обмежують застосування цих методів, є їх порівняно вузька спрямованість, відсутність доступного і якісного програмного забезпечення для практичних застосувань, необхідність спеціальних знань для складання моделей. При цьому потрібно зазначити, що на практиці, наприклад, метод імітації відпалювання часто працює швидше і дозволяє знайти краще рішення, ніж генетичний алгоритм. Тому доцільні додаткові дослідження щодо їх застосування для рішення різних економічних задач.

Імітаційне моделювання.

Імітаційна модель являє собою програмну реалізацію реальної, або гіпотетичної системи, що складається з деякого набору пов'язаних між собою об'єктів з встановленими властивостями. В процесі імітаційного експерименту відбувається зміна характеристик об'єктів моделі у часі, що взаємно впливає на пов'язані об'єкти. Це веде до зміни стану всієї моделі.

Існує кілька різновидів імітаційного моделювання [122, 231]:

– *Моделювання системної динаміки.* Виникнення цього напрямку пов'язане з ім'ям Дж. Форрестера, який в середині 1950-х років розробив його основні засади. Системна динаміка передбачає найвищий рівень агрегування компонентів з усіх методів імітаційного моделювання. Даний метод заснований на моделюванні руху потоків будь-якої природи в системі. Він дозволяє враховувати затримки потоків будь-якого порядку, петлі зворотного зв'язку, відображати процеси накопичення і переміщення ресурсів, динамічно регулювати інтенсивність потоків. Завдяки ряду спрощень, прийнятих в моделях системної динаміки (абстрагування від індивідуальних характеристик об'єктів і фізичних характеристик навколишнього середовища, безперервність всіх змінних і процесів), вони дозволяють простими засобами отримати адекватний опис процесів в досить складних системах.

Модель системної динаміки може служити для аналізу і розуміння причинно-наслідкових зв'язків між будь-якими її компонентами. Дозволяє порівняти варіанти рішень з управління системою. Метод є універсальним і застосовується для моделювання найрізноманітніших систем, від бізнес-процесів і моделей виробництва до моделей розвитку епідемій в медичних дослідженнях.

– *Дискретно-подійне моделювання* передбачає розгляд функціонування системи в часі і аналіз впливу на її стан зовнішніх подій, які подаються у вигляді «заявок». Найбільш відомим прикладом застосування цього різновиду імітаційного моделювання є аналіз систем масового обслуговування.

Сферою застосування дискретно-подійного моделювання можуть служити будь-які системи, пов'язані з обслуговуванням потоку об'єктів – системи передачі інформації, логістичні, транспортні, виробничі системи і

багато інших.

– *Агентне моделювання*. Передбачає визначення поведінки одиначної простої структури – агента – у взаємодії з іншими такими ж агентами і навколишнім середовищем. Агент може розглядатися як деяка сутність, яка має активність, автономну поведінку, може приймати рішення відповідно до деякого набору правил, може взаємодіяти з оточенням і іншими агентами, а також може еволюціонувати. Вивчення поведінки сукупності агентів дозволяє отримати інформацію про властивості системи, що вивчається, на макрорівні.

Агентний підхід доцільно застосовувати в тому випадку, коли індивідуальна поведінка об'єктів має великий вплив на поведінку системи в цілому. До таких завдань відносяться моделювання ринків, конкуренції, динаміки населення і інших. Крім того, агентний підхід може застосовуватися і спільно з іншими різновидами імітаційного моделювання, зокрема, в рамках моделей системної динаміки.

Взаємозв'язок між методами інтелектуальних обчислень в рамках виділених вище підходів показано на рис. 1.10.

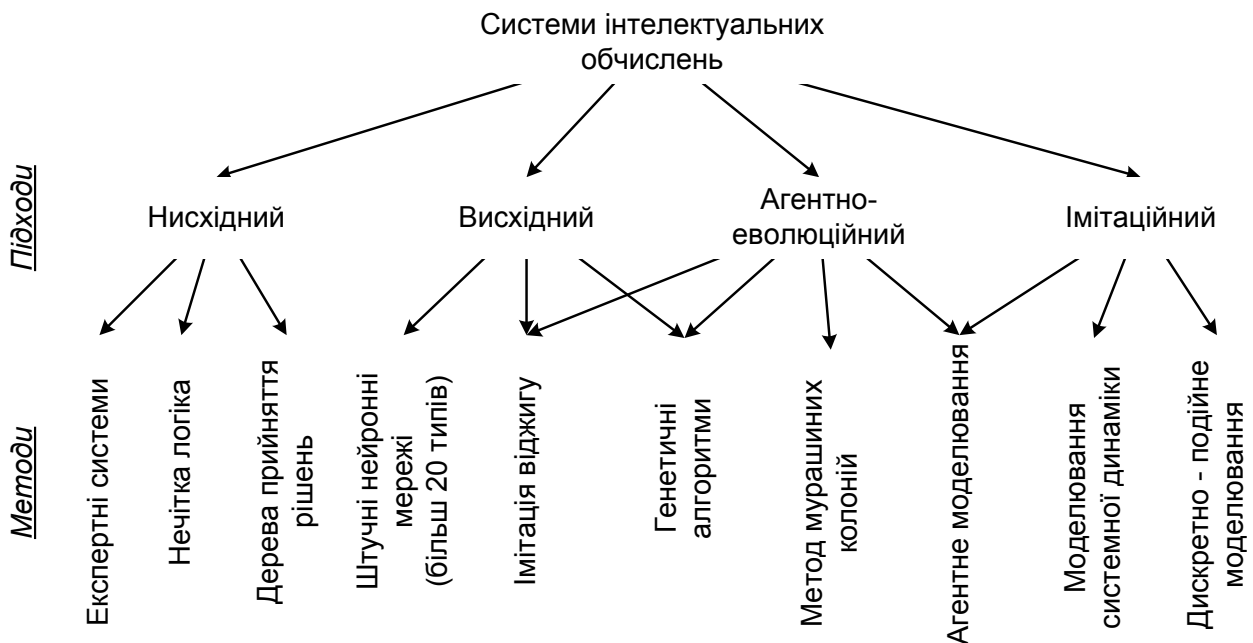


Рисунок 1.10 – Класифікація методів інтелектуальних обчислень

Як видно з розгляду класифікації, показаної на рис. 1.10, деякі методи інтелектуальних обчислень можуть бути розглянуті в рамках декількох підходів. Це не є наслідком недосконалості класифікації, а показує багатогранність розглянутих концепцій.

При цьому слід мати на увазі, що універсальні методи в рішенні деяких задач програють по ефективності вузькоспеціалізованим методам. Наприклад, створений в 1992 році метод опорних векторів дозволяє вирішувати задачі класифікації з ефективністю, еквівалентною ШНМ тих часів, але вимагає набагато менше часу на навчання [101]. Спеціально розроблені алгоритми рішення задачі комівояжера в даний час дозволяють знайти абсолютно краще рішення за менший час, ніж ГА знайде приблизне. Проте витрати на розробку і

реалізацію таких спеціальних методів можуть виявитися набагато більше отриманого ефекту.

Отже, інтелектуальні обчислення можуть бути використані для вирішення більшості економічних задач, забезпечуючи якщо не краще, то принаймні достатньо добре рішення. В даний час основним напрямком досліджень є вдосконалення самих інтелектуальних методів з метою забезпечити кращу якість, або кращу швидкість пошуку рішень. Однак одержуваний приріст ефективності, як показують проведені дослідження, виявляється порівняно невеликим [63]. Нижче розглядається альтернативний шлях – розробка методології пошуку комбінації методів, які забезпечують отримання найбільш ефективного рішення.

1.3 Методологічні підходи до синтезу інноваційних інтелектуальних систем прийняття рішень

Визначимо категорій «методологія» і «методологічний підхід» в тому значенні, в якому вони будуть використовуватися надалі. Необхідність для цього обумовлена відсутністю в даний час єдиної думки щодо їх трактування.

Для методології довгий час використовувалося визначення її як «вчення про методи діяльності», яке дав ще Р. Декарт [105]. Пізніше з'являються вузьке і широке визначення методології. Наприклад, у філософському словнику 1972 року стверджується, що це «1) сукупність прийомів дослідження, що застосовуються в будь-якій науці; 2) вчення про метод пізнання і перетворення світу» [225].

У 1970-х роках в працях Г. П. Щедровицького, Є. Г. Юдіна, В. І. Садовського та інших набув поширення принцип, відповідно до якого методологія науки стала розглядатися на чотирьох рівнях [238, 242]:

- філософському;
- загальнонауковому;
- конкретно-науковому;
- технологічному.

Така класифікація дозволила для кожного рівня розробити свої прийоми, підвищення ефективності досліджень, що зробило її загальновизнаною. Так, підходи до синтезу інтелектуальних систем прийняття рішень, що розглядаються нижче, слід віднести до останнього, технологічного рівня.

Однак огляд підходів до визначення методології був би неповним без згадування про роботи А.М.Новікова і Д.А. Новікова, метою яких стала розробка єдиної методології продуктивної, тобто спрямованої на отримання об'єктивно, або суб'єктивно нового результату, діяльності. Методологія визначена ними, як «вчення про організацію діяльності». При цьому автори виділяють як загальні методологічні принципи, так і принципи, специфічні для певних видів діяльності – наукової, практичної, навчальної, ігровий [181]. Вагомим результатом згаданого дослідження є запропонована авторами схема «структури методології», тобто рекомендованої послідовності її викладу, яка

включає [181]:

1. Характеристики діяльності (особливості, принципи, умови, норми)
2. Логічну структуру діяльності (суб'єкт, об'єкт, предмет, форми, засоби, методи, результат)
3. Часову структуру діяльності (фази, стадії, етапи).

Саме цих принципів визначення категорій методології будемо дотримуватися надалі.

Методологічною основою пропонованих підходів до синтезу інтелектуальних систем прийняття рішень є теорія систем автоматизованого проектування. Будучи порівняно молодою, ця область науки встигла акумулювати значний обсяг успішних розробок в області синтезу не тільки інженерно-механічних, але також електронних та інформаційних систем, які за структурою вхідних і вихідних інформаційних потоків а також за методами їх перетворення близькі до систем прийняття рішень [182, 85].

Відповідно до визначення, даного в [90], синтез являє собою проектну процедуру, метою якої є з'єднання різних елементів, властивостей, сторін і тому подібних характеристик об'єкта в єдине ціле, систему. В результаті синтезу виникає синергетичний ефект, тобто система отримує нові якості відносно своїх елементів.

Основними методологічними підходами до синтезу складних інформаційних систем, до яких відносяться ІСПР, є *структурний* і *параметричний*.

В рамках *структурного* підходу розглядаються методи вибору і оптимізації структури інформаційної системи. Відзначається, що саме структура несе в собі основну інформацію про функціональне призначення об'єкта і визначає його основні характеристики [89].

Для ІСПР поняття структури може розглядатися на мезорівні і мікрорівні.

Мезорівень становить набір використаних моделей, методів аналізу і обробки даних, генерації альтернатив і прийняття рішень, послідовність їх виконання, систему зв'язків між елементами.

На мікрорівні задача структурного синтезу вирішується для окремих елементів ІСПР. Це особливо актуально при використанні інтелектуальних методів пошуку рішень, оскільки для більшості з них вибір структури сильно впливає на результат.

Твердження 1.1.

Якщо задачу, що вирішується, можна віднести до категорії типових, то вибір структури може бути в достатній мірі формалізований і зведений до адаптації стандартних структур під характеристики конкретної задачі.

Доведення твердження 1.1.

Відповідно до даного вище визначення, структура несе основну інформацію про функціональне призначення об'єкта та його характеристики. Відтак, об'єкти зі схожим призначенням і характеристиками можуть мати схожу структуру. Задачі в рамках одного типу *a priori* мають схоже призначення і характеристики. Отже, для їх вирішення можуть використовуватися ІСПР зі схожими структурами.

Твердження 1.1 доведено.

Прикладом формалізації процедур вибору структури є вибір кількості нейронів у вхідному і вихідному шарах нейронної мережі. Для задач із однаковою розмірністю даних оптимальні структури у загальному випадку також будуть співпадати..

Якщо задача не є типовою, застосовується ітеративний синтез, при якому задаються загальні вимоги до ІСПР і основні обмеження (наприклад – за швидкістю розробки, швидкодією, точністю, ресурсами), після чого в кілька ітерацій формується відповідна їм система. Для багатьох інтелектуальних методів пошуку рішень ітеративний процес синтезу є обов'язковим, принаймні, на мікрорівні. Це відноситься, зокрема до визначення кількості нейронів в прихованих шарах ШНМ [232, 180], визначення розмірів популяції та кількості поколінь в генетичних алгоритмах [177, 129], рівня розгалуження в деревах рішень [5], а також деяким іншим задачам.

При *параметричному* підході синтез рішення задачі зводиться до оптимізації параметрів моделі з заданою структурою. Фактично, при цьому відбувається налаштування моделі на задані умови зовнішнього середовища. Параметричний синтез дозволяє досягти високої ефективності за умов, близьких до заданих. Але з іншого боку, при зміні зовнішніх умов, ефективність отриманої системи різко знижується.

Залежно від постановки, задача параметричного синтезу може розглядатися, як пошук в N -вимірному просторі (де N – кількість параметрів) точки, для якої умови працездатності системи або просто виконуються, або виконуються найкращим чином. В [90] показано, що в першому випадку задача параметричного синтезу зводиться до задачі оптимізації, а в другому – еквівалентна їй.

У процесі синтезу складної інформаційної системи, якою є ІСПР, використовуються обидва підходи. Структурний необхідний для визначення оптимального набору методів вирішення поставлених завдань, і їх взаємозв'язку. Параметричний – для оптимізації елементів системи. В такому випадку слід вести мову про структурно-параметричний синтез ІСПР [104].

Аналіз джерел, присвячених проблемам структурно-параметричного синтезу [104, 99, 142, 74] дозволяє в узагальненому вигляді представити основні його етапи наступним чином (рис. 1.11)

На першому етапі відбувається ідентифікація приналежності економічної задачі до одного, або декількох типів. Це дозволяє обмежити на наступному кроці кількість базових моделей, що використовуються в процедурі синтезу структури. Процедура синтезу є ітеративною. Оптимальна структура вибирається відповідно до заданих критеріїв ефективності, ризикованості, надійності і тому подібними. Після формування структури ІСПР здійснюється оптимізація параметрів системи і її налаштування на задачу, що вирішується (навчання). Для інтелектуальних методів пошуку рішень ця процедура в більшості випадків також ітеративна.

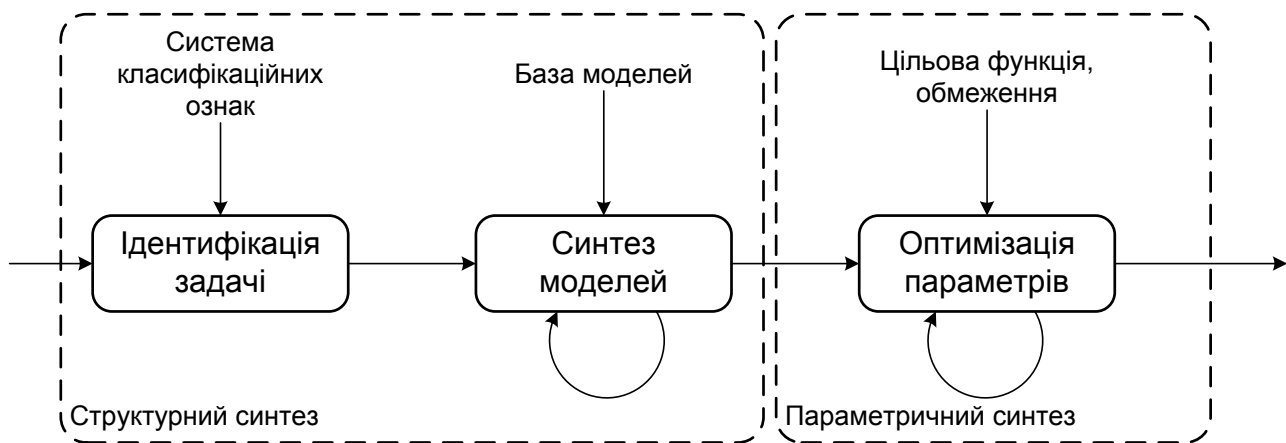


Рисунок 1.11 – Основні етапи структурно-параметричного синтезу

Зупинимося на методах вирішення завдань структурного синтезу ІСПР. З погляду можливостей формалізації задача структурного синтезу відноситься до найбільш складних, оскільки є NP-повною [74]. В даний час не існує універсального методу вирішення таких задач. Згідно [90], для структурного синтезу використовуються наступні методи:

формальний – синтез проводиться шляхом механічного перекладу постановки задачі до структури моделі за заздалегідь заданими загальними правилами. Не вимагає участі людини в основних етапах синтезу;

евристичний – є протилежністю попередньому. Основні етапи синтезу здійснюються людиною з використанням його знань, навичок і інтуїції;

комбінаторний – заснований на переборі різних комбінацій в просторі рішень серед аналогічних, або схожих задач;

спеціалізовані методи – розробляються під конкретні задачі в тому випадку, якщо жоден з перелічених методів не підходить.

Найбільший інтерес з позицій даного дослідження представляють комбінаторні методи синтезу. Їх використання з одного боку потенційно дозволяє досягти високого рівня автоматизації процесів синтезу ІСПР, а з іншого – дозволяє забезпечити достатню гнучкість одержуваних рішень за рахунок участі людини в цьому процесі. В рамках комбінаторного синтезу виділяють наступні методи [237]:

- морфологічний синтез;
- синтез по альтернативним деревам;
- синтез по багатодольним графам;
- синтез по орієнтованим гіперграфам;
- синтез на основі мереж Петрі.

Всі ці методи є в достатній мірі універсальними і можуть бути застосовані, якщо для об'єкта моделювання виконано наступні умови [237]:

1. Об'єкт має структуру;
2. Об'єкт належить до певного класу, що складається з великої кількості об'єктів, що мають однакове функціональне призначення;
3. Складові частини об'єктів можуть добре комбінуватися один з одним.

Твердження 1.2.

Для синтезу інтелектуальних систем прийняття рішень можна

використовувати комбінаторні методи.

Доведення твердження 1.2.

Для ІСПР умова 1 виконується *a priori*, так як структуру ІСПР утворюють формують її моделі та методи.

Умова 2 є більш сильною. Воно означає, що при синтезі об'єкта необхідно спиратися на досить велику базу моделей прийняття рішень, побудованих з використанням тих же складових частин методів. Однак при вирішенні економічних задач, кількість методів порівняно невелика, що знижує вимога до розміру бази моделей. Крім того, база може бути наповнена за рахунок синтетичних моделей, отриманих на умовних прикладах.

Умова 3 для інтелектуальних методів прийняття рішень в загальному випадку не виконується. Розглянуті методи різні не тільки за представленням інформації на входах і виходах (ряди, матриці, графи, дискретні і безперервні, чотки та нечіткі дані і т.п.), але і по їх адекватності для вирішення різних економічних задач. Тому, для того щоб забезпечити можливість комбінаторного синтезу системи з таких компонентів нижче будуть розглянуті спеціальні методи, орієнтовані на елементи з обмеженою сполучуваністю [90]. Використання таких методів забезпечує виконання умови 3.

Твердження 1.2 доведено.

Аналізуючи далі методи комбінаторного синтезу, слід зазначити, що всі вони, за винятком синтезу на основі мереж Петрі, можуть бути розглянуті, як розширення і уточнення методу морфологічного синтезу, запропонованого Ф. Цвіккі [71].

У класичному варіанті методу морфологічного синтезу виділяється чотири етапи. Пояснимо їх сутність, стосовно синтезу ІСПР.

1. З'ясовується мета задачі. У випадку, що розглядається, метою є побудова структурної схеми ІСПР.

2. Виділяються *вузлові точки*, що характеризують систему з позицій раніше поставленої мети. В якості вузлових точок можуть розглядатися окремі частини розв'язуваної задачі, наприклад: розвідувальний аналіз даних, аналіз даних, обробка даних, оптимізація. Набір вузлових точок може змінюватися в залежності від ідентифікованого типу задачі (див. рис. 1.11). Кількість вузлових точок в загальному випадку не лімітується, але для прискорення синтезу рекомендується починати з визначення порівняно невеликої кількості головних вузлів, а потім доповнювати модель другорядними вузлами, як наприклад інтерфейс користувача, програмне середовище реалізації і тому подібні.

3. Для кожної вузлової точки формується набір *елементів*, що представляють собою варіанти її реалізації – методи, які можуть бути використані для вирішення даної частини задачі. Перелік повинен охоплювати всі можливі варіанти, включаючи ті, які спочатку можуть здатися недостатньо ефективними, наприклад, традиційні методи аналізу даних, або моделі лінійної оптимізації.

4. Проводиться перебір всіх можливих варіантів поєднань елементів. Варіанти, які містять несумісні елементи відсіваються. Решта варіантів перевіряються на достовірність, і відповідність умовам задачі.

Результатом морфологічного синтезу є деякий набір потенційно-реалізованих варіантів, який може бути додатково проаналізовано з метою вибору найкращого.

У класичному вигляді задача морфологічного синтезу швидко ускладнюється зі збільшенням кількості вузлів і елементів в них. Дійсно, нехай n - кількість вузлів в системі, а $e_1, e_2, \dots, e_n$ – кількість елементів, розташованих у відповідних вузлах. Тоді загальна кількість варіантів синтезу можна знайти за формулою:

$$v = \prod_{i=1}^n e_i. \quad (1.3)$$

З (1.3) є очевидним, що складність задачі синтезу зростає в геометричній прогресії. Так якщо кожен вузол буде містити всього 6 варіантів (що відповідає кількості основних інтелектуальних засобів пошуку рішень, описаних в п. 1.2, без урахування їх варіацій), то при чотирьох вузлах знадобиться проаналізувати 1296 варіантів структури, а при 7 вже майже 280 тисяч. Якщо ж крім базових методів розглядати ще і їх різновиди, кількість яких для деяких методів вимірюється десятками, то загальна кількість варіантів збільшиться на кілька порядків.

Тому розвиток методу морфологічного синтезу здійснювався, перш за все, в напрямку скорочення розмірності множини варіантів: за рахунок відкидання явно неможливих, або неефективних варіантів поєднань елементів вже на етапі синтезу. В даний час отримали розвиток дві основні концепції формування вихідної множини альтернатив, кожна з яких має свої переваги, недоліки і сферу використання.

Перша концепція передбачає представлення альтернатив в просторі класифікаційних ознак та їх значень. При цьому формується морфологічна множина, в якій присутні лише допустимі рішення. Вибір альтернативи здійснюється шляхом порівняння значень класифікаційних ознак розв'язуваної задачі з елементами морфологічної множини. Використання цієї концепції доцільно при доступності достатнього обсягу даних, на яких відбувається навчання.

Друга концепція передбачає представлення альтернатив в просторі елементів і зв'язків між ними. Множина альтернатив в цьому випадку представляє собою граф, що фактично є універсальною множиною, тобто містить всі допустимі рішення [75]. Пошук структурних рішень проводиться методом занурення розв'язуваної задачі в універсальну множину [218]. Перевагами цієї концепції є компактне і більш наочне уявлення даних, можливість роботи без попереднього формування великої бази моделей, кращі можливості щодо інтерактивного синтезу. Недоліком є складність формалізації критеріїв вибору.

Обидві концепції представлення альтернатив є взаємно-еквівалентними, тобто від уявлення альтернатив в просторі класифікаційних ознак і їх значень

можна перейти до представлення в просторі елементів і зв'язків між ними і навпаки. Це дозволяє вибирати зручний варіант, в залежності від умов, що склалися. Покажемо, як здійснюється перехід від першої концепції представлення альтернатив до другої.

Нехай морфологічна таблиця, що задає простір альтернатив виглядає наступним чином (табл. 1.1):

Таблиця 1.1 – Морфологічна таблиця синтезу структури умовної системи

Вузлова точка	Елементи реалізації функції вузловий точки			
n1	e ₁	e ₂	e ₃	e ₄
n2	e ₅	e ₆	e ₇	—
n3	e ₈	e ₉	e ₁₀	e ₁₁
n4	e ₁₂	e ₁₃	—	—

Скористуємося способом завдання графа за допомогою матриці інцидентності, в якій вказуються зв'язки між вершинами графа (рядки матриці) і його ребрами (стовпці). Нульове значення вказує на відсутність зв'язку ребра і вершини, а 1 - на її наявність [77]. Аналогічним чином може бути задана інцидентність елементів реалізації вузлових точок.

Нехай n – кількість вузлів, а e – кількість елементів в морфологічній таблиці. Матриця інцидентності в цьому випадку буде мати n рядків і e стовпців. На приналежність елемента e_i до вузлової точки n_j вказує "1" що стоїть на перетині j строки та i стовпця (табл. 1.2).

Таблиця 1.2 – Матриця інцидентності елементів реалізації вузлових точках

	e1	e2	e3	e4	e5	e6	e7	e8	e9	e10	e11	e12	e13
n1	1	1	1	1	0	0	0	0	0	0	0	0	0
n2	0	0	0	0	1	1	1	0	0	0	0	0	0
n3	0	0	0	0	0	0	0	1	1	1	1	0	0
n4	0	0	0	0	0	0	0	0	0	0	0	1	1

Подальший синтез системи зводиться до вибору n елементів таких, що в кожному рядку матриці інцидентності має бути розташований один і тільки один вибраний елемент, інцидентний відповідному вузлу.

Отриманий набір елементів, упорядкований за номерами, являє собою варіант структури синтезованої системи, де кожна вузлова точка реалізується одним з можливих елементів. Якщо для будь-якого виду можливо аналітичними методами визначити рівень ефективності, подальший синтез структури зводиться до вирішення задачі оптимізації.

Матриця інцидентності (табл. 1.2), яка представлена в графічному вигляді, утворює повний n -дольний граф, показаний на рис. 1.12. Будь-яка структура, що синтезована на його основі, буде являти собою підграф із

зв'язними долями [209].

Недоліком такого методу синтезу системи є надмірність сполучень елементів, що ускладнює і уповільнює вибір її оптимальної структури. Якщо є можливість заздалегідь визначити явно нежиттєздатні варіанти структур, доцільно забезпечити синтез структури з урахуванням допустимих поєднань елементів.

Виділяють наступні види умов, що визначають можливі поєднання елементів x і y [90]:

1. *Примушування*. Вибір x тягне за собою вибір y .
2. *Необхідність*. Для вибору y потрібно вибрати x .
3. *Бінарна заборона на поєднання*. Елементи x і y не можуть входити до одного рішення.
4. *Подвійне примушення*. Елементи x і y повинні входити в рішення одночасно.

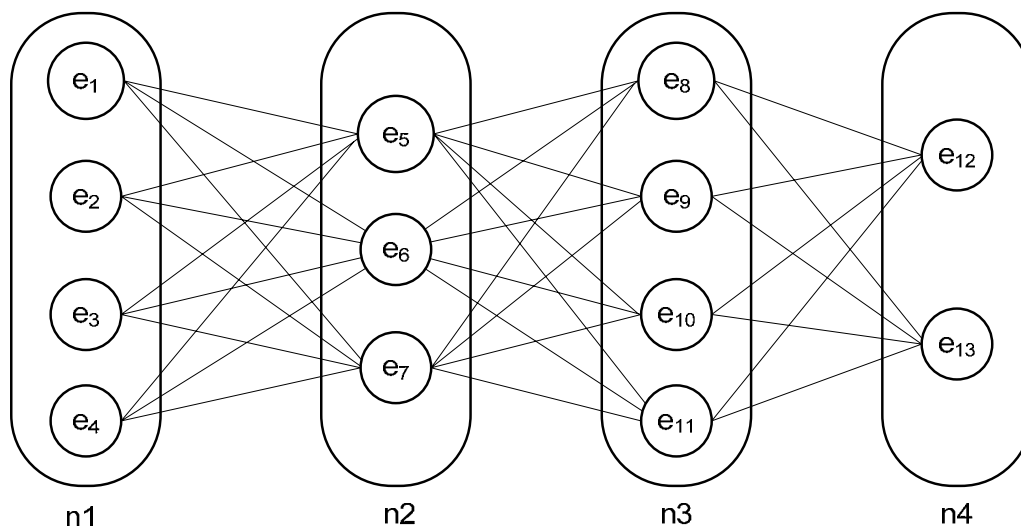


Рисунок 1.12 – Простір синтезу системи, представлене у вигляді n -дольного графа

Серед перелічених умов найбільш часто зустрічається бінарна заборона на поєднання (3). Стосовно задачі синтезу ІСПР це означає неможливість використання методів y для вирішення задачі x . Очевидно, що в просторі рішень, представленому у вигляді графа, задача визначення бінарних заборон еквівалентна задачі визначення допустимих зв'язків між елементами. У той же час остання краще відповідає логіці вербального опису.

Сформулюємо задачу синтезу ІСПР з урахуванням заздалегідь заданих допустимих зв'язків між елементами.

Для формалізації умов скористаємося термінологією і математичним апаратом аналізу гіперграфів, тобто таких графів, в яких кожне ребро може з'єднувати більше двох вершин.

Введемо такі визначення:

Визначення 1.7.

Модуль – група методів, використовуваних для вирішення задач певного

класу. Наприклад, розвідувальний аналіз, аналіз даних, обробка даних, оптимізація, прийняття рішення. Для кожного модуля будується свій гіперграф.

Визначення 1.8.

Кластер – набір умов, що характеризують предметну область за певною ознакою. Наприклад, представлення даних, обсяг даних, різновид задачі.

Визначення 1.9.

Умова – складова частина кластера, що представляє собою варіант параметра, що характеризує предметну область. Наприклад, кластеру «представлення даних» відповідає множина умов {чітке, нечітке}. Умови є *ребрами* гіперграфу.

Визначення 1.10.

Методи – способи вирішення задачі, які припустимі в рамках даного модуля. Методи є *вершинами* гіперграфу. Умовою припустимості методу є інцидентність йому хоча б одного ребра з кожного кластера.

Твердження 1.3.

При синтезі структури ІСПР задача розбивається на основні етапи – модулі, в рамках яких необхідно знайти допустимі набори методів, придатних для її вирішення.

Доведення твердження 1.3.

Пояснимо пропонований підхід до синтезу структури ІСПР.

Позначимо множину кластерів – L , множину умов – C , множину методів – M . При цьому очевидно, що кількість умов має бути не менше кількості виділених кластерів $|C| \geq |L|$.

Хід синтезу структури ІСПР розглянемо на прикладі. Припустимо, в рамках синтезу структури необхідно знайти допустимі методи для модуля «аналіз даних», в який входять такі методи, як дерева прийняття рішень, експертні системи і штучні нейронні мережі.

Поле параметрів синтезу задано у вигляді трьох кластерів:

представлення даних {чітке, нечітке};

кількість даних {мало, середньо, багато};

стаціонарність даних {висока, середня, низька}.

Встановимо такі бієктивні відображення:

Для множини вершин графа:

{Дерева прийняття рішень, експертні системи, ШНМ} = { m_1 , m_2 , m_3 }.

Для кластера «представлення даних»:

{Чітке, нечітке} = { c_1 , c_2 }.

Для кластера «кількість даних»:

{Мало, середньо, багато} = { c_3 , c_4 , c_5 }.

Для кластера «стаціонарність даних»:

{Висока, середня, низька} = { c_6 , c_7 , c_8 }.

Для модуля «аналіз даних»:

{Представлення даних, кількість даних, стаціонарність даних} = { l_1 , l_2 , l_3 }.

Простір допустимих альтернатив може бути заданий або через набір

множин, або через матрицю інцидентності. Розглянемо обидва варіанти.

Через набір множин простір допустимих альтернатив описуються наступними виразами:

$$I1 \begin{cases} c1 \{m1, m2, m3\}; \\ c2 \{m2, m3\}. \end{cases} \quad (1.4)$$

$$I2 \begin{cases} c3 \{m1, m2\}; \\ c4 \{m1, m2, m3\}; \\ c5 \{m2, m3\}. \end{cases} \quad (1.5)$$

$$I3 \begin{cases} c6 \{m1, m2\}; \\ c7 \{m1, m3\}; \\ c8 \{m3\}. \end{cases} \quad (1.6)$$

Матриця інцидентності, що відповідає виразам (1.4) – (1.6) має наступний вигляд (табл. 1.3).

Таблиця 1.3 – Матриця інцидентності ребер і вершин гіперграфу, що описує простір допустимих альтернатив.

	I1		I2			I3		
	c1	c2	c3	c4	c5	c6	c7	c8
m1	1	0	1	1	0	1	1	0
m2	1	1	1	1	1	1	0	0
m3	1	1	0	1	1	0	1	1

Легко показати, що між набором виразів (1.4) – (1.6) і матрицею інцидентності (табл. 1.3) існує взаємно-однозначна відповідність, що забезпечує можливість переходу від однієї форми представлення до іншої. При цьому набір множин зручніше для первинного введення умов, а матриця інцидентності – для подальшого аналізу.

При синтезі ІСПР будується підграф, множина ребер якого відповідає умовам задачі і є підмножиною ребер гіперграфу простору допустимих альтернатив.

Припустимо, в рамках ІСПР необхідно аналізувати невелику вибірку, що містить нечіткі дані. Умови зовнішнього середовища є достатньо стабільними, що забезпечує високу стаціонарність даних. Таким чином, умовам задачі відповідає підграф, що складається з ребер c2, c3, c6. На підставі табл. 1.3. побудуємо матрицю інцидентності цього графа (табл. 1.4).

Припустимі рішення задачі утворює такий набір вершин, який є *трансверсаллю* гіперграфу, тобто має непорожній перетин з кожним ребром. З табл. 1.4. випливає, що задача має єдине рішення – вершину m2, що означає відповідність її умовам тільки інструментарію експертних систем.

Таблиця 1.4 – Матриця інцидентності підграфа, що відповідає умовам задачі

	c2	c3	c6
m1	0	1	1
m2	1	1	1
m3	1	0	0

Апарат теорії графів також дозволяє обробляти ситуації, коли допустиме рішення не знайдено, або не задовольняє додатковим критеріям, які не врахованим при формалізації умов.

Припустимо, що попередній аналіз показує квазістаціонарність вхідних даних, тобто характеристики їх розподілу мають тенденцію до поступової зміни. В цьому випадку матриця інцидентності підграфа рішення прийме наступний вигляд (табл. 1.5)

Таблиця 1.5 – Матриця інцидентності підграфа, що відповідає умовам задачі при квазістаціонарності вхідних даних

	c2	c3	c7
m1	0	1	1
m2	1	1	0
m3	1	0	1

Як видно з табл. 1.5, жоден з методів не забезпечує виконання всіх умов задачі. Але при цьому аналіз матриці інцидентності дозволяє сформулювати вимоги, при виконанні яких стане можливим використання того, чи іншого методу. Так, з аналізу табл. 1.5. видно, що для використання дерев рішень необхідно перейти від нечіткого представлення даних до традиційного. Для використання експертних систем необхідно забезпечити стаціонарність даних в період роботи системи (тобто потрібне періодичне переналадження параметрів системи). Для використання ШНМ необхідно забезпечити доступ до більшої виборки даних. Подальший вибір залежить від того, які з цих вимог виявляться більш прийнятними в конкретному випадку.

Підсумовуючі розгляд пропонованої методології синтезу ІСПР необхідно відзначити наступні ключові моменти.

1. *Визначення загальної структури системи у вигляді модулів, що входять до неї.* Це визначає послідовність синтезу і набір модулів системи, що синтезується. Крім того, результати, які отримано від вхідних модулів, можуть визначати умови синтезу наступних. При цьому можуть бути запропоновані наступні варіанти для основних модулів системи: *попередній аналіз даних, обробка даних, аналіз даних, оптимізація варіантів.*

2. *Визначення набору кластерів, що входять в модулі.* Вирішення цієї задачі відповідає визначенню вузлових точок в методі морфологічного синтезу.

Для ІСПР множина можливих кластерів порівняно невелика, тому може бути в основному задана аналітично, як і набір модулів. При необхідності множина кластерів може бути легко доповнена.

3. Одним з найважливіших підготовчих етапів синтезу ІСПР є *формалізація обмежень на сполучуваність елементів*, що формують структуру моделі. Розглянутий методологічний підхід до синтезу ІСПР забезпечує можливість деталізації умов, а також оперативність доповнення і уточнення бази методів. Наприклад, замість методу «дерева рішень», що фігурує в розглянутому вище прикладі, можуть бути введені різні алгоритми синтезу цих дерев, що відрізняються здатністю до вирішення задач регресії. При цьому нові методи наслідуватимуть всі ознаки попередніх, за винятком тих, які визначають їх ключові відмінності, що істотно скорочує трудомісткість синтезу.

Твердження 1.3 доведено.

Таким чином, проведений аналіз характеристик процесу синтезу складних систем, а також логічної і часової структур, які характерні для інформаційних систем, зокрема, інтелектуальних систем прийняття рішень, дозволив запропонувати і обґрунтувати методологічний підхід, заснований на використанні методів графового представлення інформації про елементи системи що синтезується і припустимі зв'язки між ними. Особливостями запропонованого підходу є можливість урахування бінарних заборон на поєднання елементів, а також гнучкість одержуваної структурної моделі.

1.4 Концепція моделювання інноваційних інтелектуальних систем прийняття рішень в економіці

Авторська концепція є завершальною частиною індуктивної фази дослідження та є стислим формулюванням, що відображає його сутність [181].

На основі аналізу інтелектуальних систем підтримки прийняття рішень, їх класифікації і систематизації, які було проведено в 1.1 – 1.3, запропоновано концепцію моделювання інноваційних інтелектуальних систем прийняття рішень в економіці, що визначає стратегію дій в рамках наступної дедуктивної фази дослідження, тобто його розвитку від абстрактного до конкретного. Саме на підставі концепції виділяється комбінація моделей і методів, що забезпечують досягнення мети дослідження.

Розглянемо принципи вибору методів вирішення економічних задач, з позицій системного аналізу та теорії прийняття рішень. Для цього скористаємося стандартними умовними позначеннями, які приводяться, наприклад, в [81]:

Ω – множина можливих варіантів рішень;

ω – рішення, що обране;

C – правила вибору найкращої альтернативи (задаються у вигляді функції вибору).

Тоді в загальному вигляді задача вибору рішення запишеться в такий спосіб:

$$\omega = C(\Omega). \quad (1.7)$$

Вибір відповідних методів пошуку рішень залежить від рівня визначеності множини варіантів Ω і ступеня формалізації функції вибору C (табл. 1.6).

Можна простежити відповідність даних табл. 1.6 і рівнів прийняття рішень в піраміді Давенпорта (рис. 1.3). Так, рівню 1 відповідають задачі оптимального вибору, рівню 2 задачі вибору і пошуку оптимальних рішень, рівню 3 загальна задача прийняття рішень.

Також легко помітити, що традиційні методи пошуку рішень дозволяють адекватно вирішувати тільки *задачі оптимального вибору*, які є добре структурованими. В цьому випадку отримане рішення буде об'єктивним і найкращим в наявних умовах. Однак ускладнення економічних процесів, що спостерігається в останні десятиліття, суттєво обмежує їх сучасне застосування.

Якщо в задачі існує формальний критерій оптимальності, але характеристики простору рішень виключають застосування аналітичних і переборних методів, наприклад – задача є NP-повною, то її можна віднести до типу *задач пошуку оптимальних рішень*. Універсальним методом їх вирішення є генетичні алгоритми, однак застосовується і ряд інших інструментів інтелектуальних обчислень, серед яких ШНМ Хопфілда, Потса, зростаючі нейронні мережі, метод імітації відпалювання, метод мурашиних колоній та інші [177].

Таблиця 1.6 – Методи вирішення економічних задач

№	Ω	C	Тип задачі	Методи вирішення
1.	Однозначно визначена	Строго формалізовано	Задача оптимального вибору	Аналітичні методи; Дослідження операцій; Спеціальні методи оптимального вибору
2.	Визначена, але перевищує обчислювальні можливості системи	Строго формалізовано	Задача пошуку оптимальних рішень	Генетичні алгоритми, Рекурентні нейронні мережі Методи фізико-біологічної оптимізації
3.	Однозначно визначена	Не формалізовано	Задача вибору	Методи скорочення невизначеності: - Імітаційне моделювання; - Методи експертних оцінок; - Теорія корисності.
4.	Може доповнюватися	Не формалізовано	Загальна задача прийняття рішень	Нечітка логіка Нейронні мережі Методи логіко-лінгвістичного моделювання

Якщо простір вибору визначено, але неможливо об'єктивно сформулювати правило відбору кращої альтернативи, задача відноситься до категорії *задач вибору*. У цьому випадку критерій вибору і його результат суб'єктивно залежить від ОПР. Для вирішення задач вибору використовується імітаційне моделювання, методи експертних оцінок, теорія корисності та інші, які дозволяють скоротити невизначеність при виборі критеріїв.

Найхарактернішою задачею в управлінні складними системами є *загальна задача прийняття рішень* (ЗЗПР) [207]. Для неї характерна відсутність можливості визначення не тільки критеріїв оптимальності, але і сам простір вибору постійно видозмінюється, що веде до необхідності постійного моніторингу стану системи і вироблення коригувальних рішень. Підтримка прийняття рішень в ЗЗПР здійснюється шляхом формування проміжної множини альтернатив, з яких здійснює вибір ОПР, тобто зведенням ЗЗПР до класу задач вибору. Формування проміжної множини альтернатив може здійснюватися за допомогою нейромережових інструментів, методів нечіткої логіки, а також методів логіко-лінгвістичного моделювання.

Визначення 1.11.

Складною економічною задачею будемо називати таку, визначення гарантовано-кращого рішення якої за припустимий час неможливе внаслідок великої кількості можливих варіантів рішень, чи недостатньої формалізації правил вибору найкращої альтернативи.

Розглядаючи прийняття рішень, як процес, що відбувається в часі, слід зазначити, що в міру його розвитку необроблений масив інформації, який характеризує досліджувану систему, проходить через ряд послідовних трансформацій, що призводять в остаточному підсумку до вибору рішення, яке визначає подальший розвиток системи. У цій послідовності можна виділити процеси, пов'язані із спостереженням та моделюванням досліджуваної системи, ідентифікацією її стану, оцінкою і вибором альтернатив. Розглянемо їх.

1. Процес спостереження [81].

Введемо наступні позначення:

Z – простір станів аналізованої системи. Повний набір характеристик всіх елементів, що складають систему разом із середовищем її функціонування і їх можливих значень. Оскільки навіть для невеликих систем складання такого набору не тільки недоцільно, але часто і неможливо, простір Z можна розглядати лише як математичну абстракцію;

Y – множина характеристик системи, які спостерігаються. Набір характеристик, які входять в простір Z , та найкращим чином відображають стан системи і зовнішнього середовища з певної точки зору. Дані характеристики повинні бути здатними до спостереження, тобто повинен існувати спосіб визначення їх дійсних значень. Серед параметрів, що складають множину Y можна виділити вхідні Y_{in} та вихідні Y_{out} характеристики системи, прийняті рішення Y_D , їх ефективність Y_{DE} та їм подібні.

Задача спостереження, таким чином, включає відбір характеристик, що складають множину Y , відстеження їх значень і збереження отриманої інформації в базі даних. В термінах системного аналізу рішення цієї задачі

зводиться до відшукування такого відображення

$$g^{-1}: Y \rightarrow Z,$$

яке для кожної реалізації характеристик, що спостерігаються, Y , ставить в однозначну відповідність внутрішній стан об'єкта управління.

У процесі спостереження крім організації збору інформації можуть виникати задачі обробки даних, сутність і класифікацію яких докладно розглянуто вище. Множина Y також часто називається вхідною вибіркою даних. Причому вхідні вибірка може містити як чіткі, так і нечіткі дані.

2. Процес моделювання [126].

Оскільки потужність множини Y і розмірність відповідного простору станів можуть бути досить великими, задача прямого аналізу всіх можливих станів безпосередньо на підставі множини Y для скільки-небудь складної системи є трансобчислювальною. Для скорочення складності задачі доцільно звільнити її від другорядних деталей і зв'язків, тобто побудувати *модель*. Складність її з одного боку повинна бути достатньою для подальшого аналізу, а з іншого не повинна перевищувати обчислювальних можливостей системи. Слід зазначити, що терміни «модель» і «моделювання» тут розуміються в самому широкому сенсі, охоплюючи різні типи економіко-математичних моделей, вибір яких залежить від особливостей вирішуваних задач, визначеності простору рішень Ω і правил вибору альтернативи - S .

З формальної точки зору, рішення задачі моделювання може розглядатися як побудова абстрактної множини E , ізоморфної предметної області:

$$\varphi: Y \rightarrow E.$$

Основним призначенням моделі E в задачі прийняття рішень є вивчення та прогнозування реакції об'єкта на керуючі впливи.

Акцентуємо увагу на тому, що існують різні підходи до побудови економіко-математичних моделей, зокрема:

аналітичний підхід, при якому модель формується, як висновок з аналізу об'єктивних залежностей між параметрами об'єкта управління;

формальний підхід, при якому побудова моделі проводиться за допомогою обробки накопичених даних і підбору апроксимуючих (тобто приблизних) залежностей між ними;

синтезуючий підхід, при якому модель утворюється в результаті синтезу з деякого підмножини бази готових моделей.

Процес моделювання включає вирішення задач, пов'язаних з аналізом даних, синтезом структури ІСПР на макро- і мезо- рівнях, параметричному синтезі моделей, вибором інструментальних засобів їх реалізації та іншими необхідними процедурами. Таким чином, в моделюванні інноваційних ІСПР цей процес є найбільш наукомістким.

3. Процес ідентифікації [232]

Процес ідентифікації пов'язують із вирішенням задачі розпізнавання образів, тобто відшукування відповідності між характеристиками системи, які спостерігаються в даний момент (вектор S) і станами системи, які спостерігалися раніше.

Основна проблема ідентифікації полягає в тому, що внаслідок обмеження часу спостереження повний збіг значень вектора S з якимось елементом множини Y представляється неможливим (зрозуміло, якщо характеристики Y досить повно відображають стан системи, тобто за умови грамотного рішення задачі спостереження). Отже, задача ідентифікації зводиться до відшукування найбільш схожої ситуації, з тими, що спостерігалися раніше. Для цього використовуються методи, засновані на визначенні відстані між вектором S і векторами, що описують попередні стани системи в n -вимірному просторі (моделлю системи E), де n – розмірність вектора S . Найбільш схожій ситуації буде відповідати мінімальна відстань. На практиці, однак, необхідно мати на увазі різну значущість характеристик в кожному конкретному випадку.

Точність ідентифікації багато в чому визначається адекватністю моделі і, як наслідок, залежить від результатів процесу моделювання. Разом з тим достовірність ідентифікації може бути підвищена за рахунок скорочення невизначеності зовнішнього середовища шляхом оцінювання параметрів системи, які на момент ідентифікації не є достовірно відомими.

4. Процес оцінки та вибору альтернатив.

З позицій системного аналізу для загальної задачі прийняття рішень справедливі наступні твердження:

Твердження 1.4.

Не існує оптимального стану економічної системи для всіх видів впливів і станів зовнішнього середовища.

Доведення твердження 1.4.

Оскільки множина можливих станів зовнішнього середовища нічим не обмежена, її можна вважати нескінченною. У нескінченній множині, за визначенням, присутні такі стани зовнішнього середовища, для яких простори допустимих станів системи не перетинаються.

Твердження 1.4. доведено.

Слідство з твердження 1.4: Ефективність економічної системи може розглядатися тільки для конкретної мети і конкретних умов.

Твердження 1.5.

Не існує стану системи, найкращого для всіх ОПР. У схожій ситуації інша ОПР може вибрати інше рішення.

Доведення твердження 1.5.

Розглянемо ОПР, як частину зовнішнього, по відношенню до системи, середовища. Приймаючи рішення, ОПР керується не тільки об'єктивною інформацією про стан системи, але і суб'єктивними критеріями, які визначаються зокрема його знаннями, освітою, вихованням, інтенцією. Отже, різні ОПР формують різний стан зовнішнього середовища, для кожного з яких, відповідно до твердження 1.3 можуть бути оптимальним різний стан системи.

Твердження 1.5 доведено.

В інтелектуальній системі прийняття рішень роль ОПР виконує сама система, фактично моделюючи його поведінку. З урахуванням вищевикладеного зрозуміло, що в складі ІСПР необхідно мати кілька таких моделей, вибір яких здійснюється автоматично, або примусово. Кожна модель передбачає різні поєднання домінуючих критеріїв оцінки і вибору, серед яких можна виділити [236]:

- *критерії ефективності*: результативність, ресурсомісткість, оперативність;
- *критерії ризиковості*: обережність, прибутковість, компроміс;
- *критерії перспективності*: орієнтація на дальню, середню або ближню перспективу.

Розглянемо формальну постановку задачі оцінювання.

Нехай Ω - множина оцінюваних альтернатив;

i - критерій оцінки;

M_i - оцінка альтернативи по i -му критерію;

M - множина таких оцінок.

Відзначимо, що для складних економічних систем отримання оцінок M_i пов'язано із необхідністю згортки багатовимірного простору критеріїв [139]. При цьому кожній моделі поведінки економічної системи повинен відповідати свій набір коефіцієнтів важливості критеріїв. Проте, виділяти задачу згортки в окремий клас недоцільно, оскільки вона може розглядатися, як окремий випадок рішення описаної вище задачі моделювання.

Таким чином, задача оцінювання альтернатив зводиться до відшукування відображення:

$$\mu: \Omega \rightarrow M,$$

тобто такого набору функцій μ_j , які б забезпечували однозначна відповідність між альтернативою $\omega \in \Omega$ і її оцінкою відповідно до заданого критерію j . При вирішенні задачі оцінювання в такій постановці, задача вибору кращої альтернативи стає тривіальною.

Окремо слід розглянути питання генерації множини альтернатив для оцінювання Ω . Шляхи його вирішення залежать від ступеня визначеності множини Ω (см. табл. 1.6) та її скінченності. Так, якщо множина Ω визначена, скінченна, і оцінка кожної альтернативи що входить до неї не перевищує обчислювальних можливостей ІСПР, проводиться повний перебір варіантів.

Якщо множина Ω є визначеною, скінченною, але повний перебір альтернатив перевищує можливості ІСПР, використовуються методи вибіркового перебору, що виключають явно невдалі варіанти. До них відносяться генетичні алгоритми, методи імітації відпалювання та інші подібні.

Якщо множина Ω не є визначеною, то ІСПР генерує її на підставі рішень, які приймалися в минулому (підмножина Y_D) та їх комбінацій із використанням інтелектуальних методів пошуку.

Сукупність процесів спостереження, моделювання, оцінювання та вибору, що реалізує функцію пошуку і вибору рішень δ таким чином можна

представити у вигляді схеми, яка відображає концепцію даного дослідження (рис. 1.13).

Згадані вище компоненти задачі прийняття рішень – вхідна інформація z_i , правила, що регламентують прийняття рішень c та керуючі впливи ω_i в схемі на рис. 1.13 розглядаються, як компоненти відповідних множин Z , C та Ω . Наповнення множин даними є однією з задач, що потребує вирішення при створенні ІСПР і розглядається окремо.

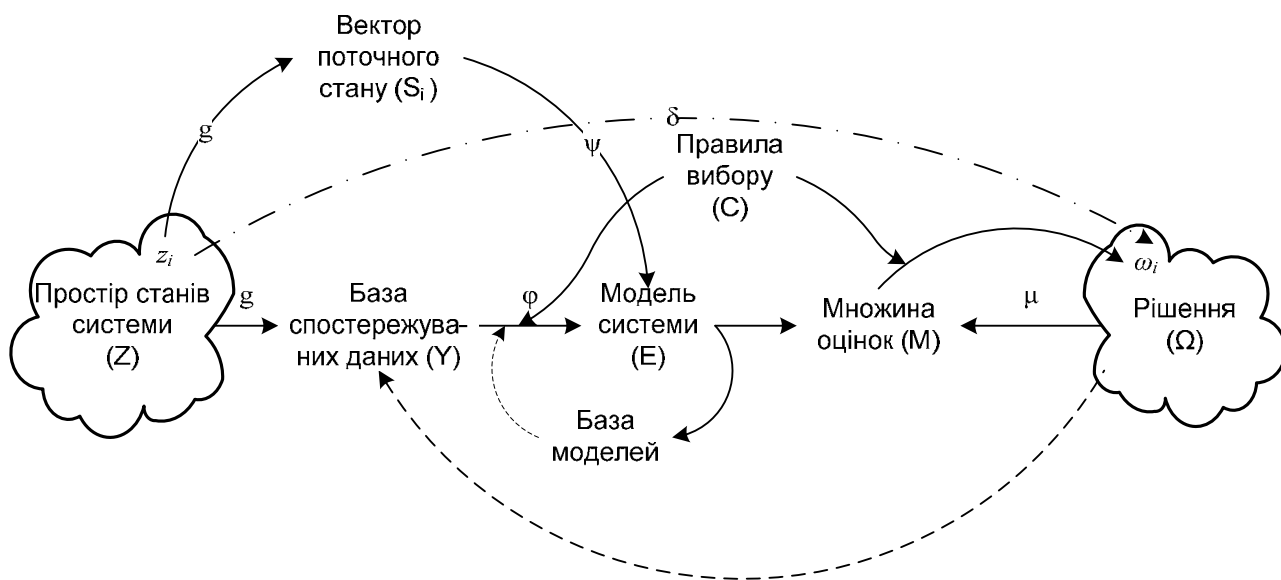


Рисунок 1.13 – Концепція моделювання інноваційної інтелектуальної системи прийняття рішень

Слід звернути увагу на те, що множина рішень Ω , що є основою для формування керуючих впливів, визначається заздалегідь, хоча може доповнюватися в процесі роботи системи. Отже, ІСПР не може згенерувати принципово нове рішення, яке не є комбінацією відомих раніше. Однак при значній кількості параметрів, що визначають рішення і чималій базі рішень це обмеження перестає бути суттєвим. Дійсно, якщо кожен варіант рішення може бути охарактеризований k параметрами, загальна кількість рішень, які можуть бути скомбіновані з n варіантів, наявних в базі, можна знайти за формулою:

$$N = n^k. \quad (1.8)$$

Так, з (1.8) випливає, що, наприклад, на підставі бази з 100 варіантів, кожен з яких описується всього 10 дискретними параметрами, може бути створено $100^{10} = 100\,000\,000\,000\,000\,000\,000\,000\,000\,000$ різних комбінацій.

Таким чином, при розробці ефективних підходів до наповнення множини Ω , методів комбінації рішень та вибору альтернатив область вирішуваних завдань залишається достатньо великою.

Основна база даних ІСПР на рис. 1.13 представлена множиною Y , склад якої більш детально розглянуто вище, при описі задачі спостереження. Множина Ω , як вже зазначалося, також може бути її компонентом, що знайшло

відображення на концептуальній схемі. Наповнення множини Y регламентується *процесом* g .

Модель системи E синтезується *процесом* φ на підставі даних вхідної вибірки (множина Y) і набору правил, що регламентують роботу ІСПР і прийняття рішень (множина C). У деяких випадках, розглянутих вище, синтез моделі прийняття рішень може проводитися з використанням бази, яка містить моделі аналогічних систем.

Структура і склад перелічених компонентів ІСПР визначають використання підходів до вирішення задачі оцінки альтернатив (*процес* μ), результатом якої є множина оцінок M . На відміну від множин Y , E , C , Ω , які було розглянуто вище, множина M не зберігається в явному вигляді, оскільки будь-яка оцінка може бути отримана з наявних даних формальними методами. Розробка цих методів і є сутністю розв'язання задачі оцінки.

Розглянемо процедуру формування рішення в ІСПР.

Поточна ситуація z_i спостерігається і відображається у вигляді вектора S_i . Його структура повторює структуру елемента множини Y в частині вхідної інформації. Далі в процесі функціонування ІСПР поточна ситуація ідентифікується за допомогою моделей E (*процес* ψ). Ідентифікація дозволяє відібрати з множини Ω , деяку підмножину рішень (пред'явлення) для подальшого вибору найкращої альтернативи.

Вибір альтернативи ω_i здійснюється на підставі результатів моделювання та їх оцінки. Вибір методів оцінювання регламентується актуальними в поточній ситуації правилами прийняття рішень (елемент множини C).

Таким чином, в розглянутій концепції ІСПР процес пошуку рішення зводиться до послідовності процесів спостереження, моделювання, ідентифікації, оцінювання та вибору. Реалізація цих процесів базується на розробках вітчизняних і зарубіжних авторів в таких областях, як статистичний аналіз, спостереження і експеримент, технології баз даних, математичне моделювання, інтелектуальних аналіз і обробка даних, методи оптимізації. Однак деякі аспекти, пов'язані з реалізацією цих процесів в ІСПР вимагають додаткового уточнення, що обумовлює необхідність вирішення низки завдань.

В рамках реалізації процесу спостереження необхідно:

- розробити генетичні моделі рішення задач обробки даних;
- розробку методів відбору даних для вирішення складних економічних задач з використанням штучних нейронних мереж;

Реалізація процесу моделювання включає:

- розробку методичного підходу до постановки задачі нейромережевого моделювання та перевірку його при вирішенні складних економічних задач;
- розробку критеріїв і методів вибору інструментальних засобів інтелектуального аналізу і обробки даних;
- дослідження ефективності сучасних інструментальних засобів інтелектуального аналізу і обробки даних.

Реалізація процесу ідентифікації вимагає:

- розробити імітаційні моделі оцінки параметрів економічних систем.

Для реалізації процесу оцінювання та вибору альтернатив слід:

- систематизувати методи оцінки ефективності вирішення складних економічних задач;
- розробити методи забезпечення порівнянності результатів при порівнянні ефективності вирішення економічних задач аналізу і обробки даних.

Розглянемо концепцію моделювання інноваційної інтелектуальної системи прийняття рішень з позицій різних методологічних рівнів за допомогою схеми, яку наведено на рис. 1.14.

Дана схема забезпечує методологічний зв'язок між об'єктом дослідження (який на процесному рівні можна представити через сукупність процесів, спостереження, моделювання, ідентифікації, оцінювання і вибору), науковими напрямками, що формують теоретико-методологічну базу дисертації, а також інструментами і методами, які використовуються для розв'язання поставлених завдань різних класів.

Основними класами задач інтелектуальних обчислень, які потрібно вирішувати при побудові інноваційної інтелектуальної системи прийняття рішень є задачі з аналізу даних, їх обробки, а також пошуку оптимальних рішень. Це не означає, що для розробки та впровадження практичної реалізації ПСПР відсутня необхідність вирішувати інші задачі, такі, як організаційне забезпечення її функціонування, моніторинг прийнятих рішень, питання організації доступу до системи, складання форм звітності та їм подібні. Але ці задачі в даний час можна вважати достатньо формалізованими. Тому вони не потребують застосування методології інтелектуальних обчислень.

На інструментальному рівні для вирішення задач з аналізу даних слід дослідити застосування таких класів методів і інструментів, як штучні нейронні мережі, нечіткі обчислення, імітаційне моделювання, формалізовані методи аналізу даних. Задачі обробки даних доцільно вирішувати за допомогою таких методів, як штучні нейронні мережі, генетичні алгоритми, нечіткі обчислення, формалізовані методи обробки даних. Для розв'язання задач оптимізації залежно від рівня визначеності множини варіантів Ω можна використовувати або формалізовані методи, або агентно-еволюційні методи пошуку оптимальних рішень.

Реалізація зазначених процесів на інструментальному рівні, відбувається з використанням розвиненого набору економіко-математичних інструментів. Серед інструментів, що мають особливе значення для ПСПР, слід виділити штучні нейронні мережі різних типів, які можуть використовуватися для аналізу і прогнозування (персептронні), для багатокритеріального порівняння альтернатив (самоорганізаційні), для розпізнавання складних образів (згорткові). У рішенні задач оптимізації важливу роль відіграє використання генетичних алгоритмів, що мають ряд унікальних переваг перед іншими методами. Крім того, для забезпечення адаптованості та самонавчання ПСПР, необхідно забезпечити динамічне оновлення інформації в базах системи, коригування моделей, критеріїв оцінювання та вибору, в залежності від результатів прийнятих і реалізованих рішень.

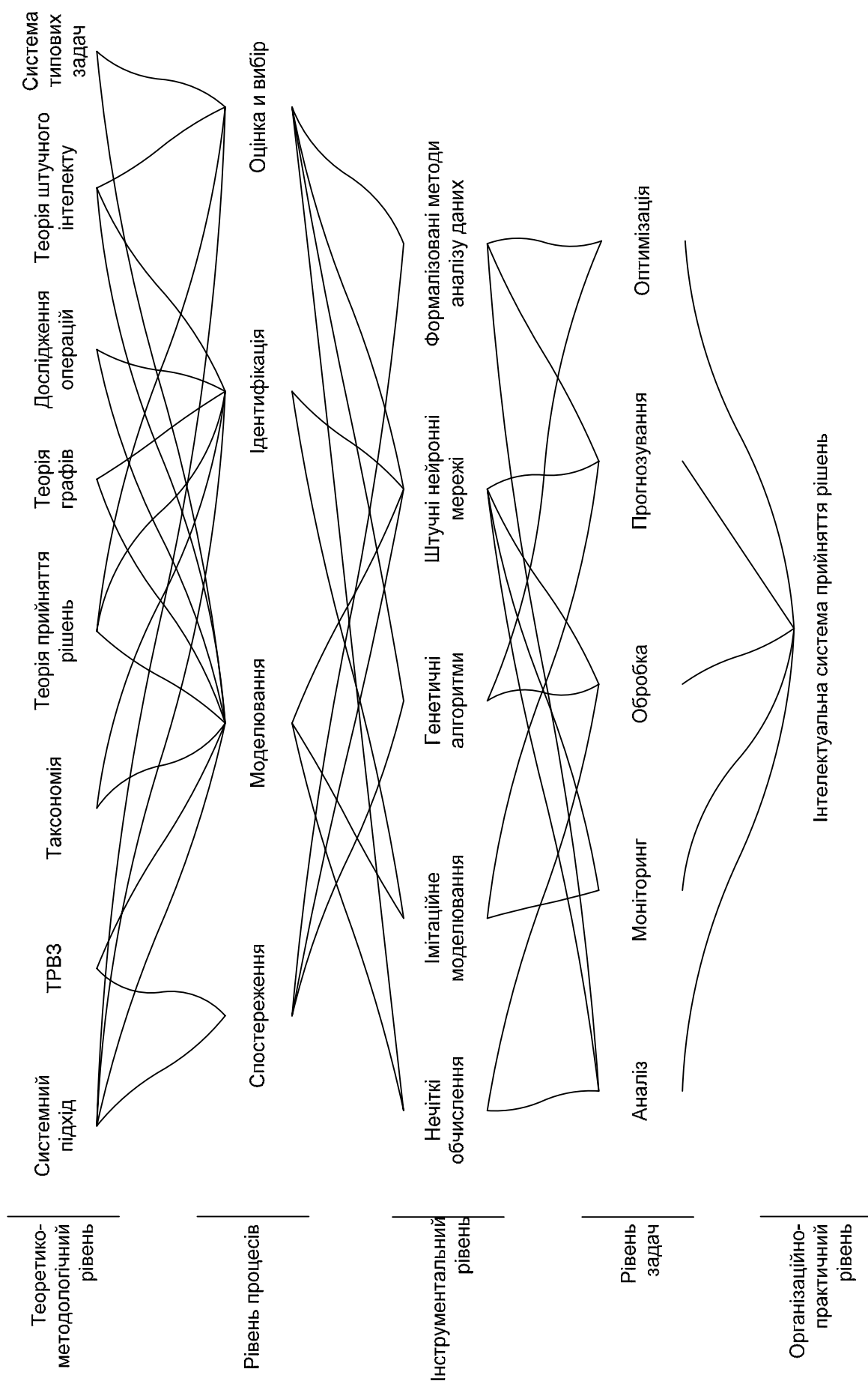


Рисунок 1.14 – Методологічна схема моделювання інтелектуальної системи прийняття рішень

Зв'язки між компонентами інструментального рівня та рівня процесів на рис. 1.14 утворюються на основі зазначених раніше аспектів, які вимагають додаткового уточнення при їх реалізації в ПСПР. Так, в рамках процесу спостереження, можуть використовуватися генетичні алгоритми, штучні нейронні мережі та формалізовані методи аналізу даних. Процес моделювання може включати використання імітаційних моделей, нечітко-логічних моделей, а також моделей на основі штучних нейронних мереж. Задачі, які поставлено в даному дослідженні на рівні процесу ідентифікації потребують використання імітаційних моделей, штучних нейронних мереж та формалізованих методів аналізу даних. Вдосконалення процесу оцінки і вибору альтернатив планується здійснювати із використанням нечітких обчислень, генетичних алгоритмів, формалізованих методів аналізу і обробки даних та штучних нейронних мереж.

Слід відмітити, що дана схема встановлює лише основні взаємозв'язки між категоріями методології моделювання інноваційних інтелектуальних систем прийняття рішень. Так, серед численних різновидів нейронних мереж існують такі, які можуть використовуватися для розв'язання задач оптимізації, але практичне їх застосування для таких випадків пов'язане із багатьма обмеженнями, що робить його недоцільним і на рис. 1.14 не показано.

Основою для розгляду процесів спостереження, моделювання, ідентифікації, оцінювання та вибору є загальновизнані теоретико-методологічні наукові напрямки, наведені на верхньому рівні методологічної схеми. Детальніше зупинимося лише на таких компонентах рівня, як теорія рішення винахідницьких задач (ТРВЗ) [78] і система типових задач.

Використання ТРВЗ на теоретико-методологічному рівні обумовлено тим, що саме в рамках цієї теорії велика увага приділяється правильній постановці задачі і доводиться, що від неї істотно залежить ефективність вирішення. В той же час доведено, що коли економічна проблема може бути зведена до однієї, або декількох базових постановок, ефективність рішень, знайдених в їх рамках, може відрізнятись [161].

Введення системи типових задач обумовлено необхідністю апріорної оцінки ефективності використовуваних інструментальних засобів і їх програмних реалізацій. Система типових задач дозволяє створити єдиний базис для такої оцінки і забезпечує порівнянність результатів.

Таким чином, в першому розділі монографії на підставі аналізу проблематики прийняття рішень в економічних системах, систематизації економічних задач та методів їх вирішення визначено методологічні підходи до синтезу інноваційних інтелектуальних систем прийняття рішень, а також розроблено концепцію їх моделювання, яка дозволяє звести задачу пошуку і прийняття рішень до набору більш простих завдань, вирішення яких може бути здійснено з використанням сучасних методів аналізу і обробки даних. Використання інтелектуальних обчислень дозволяє підвищити ефективність прийнятих рішень, скоротити витрати на підтримку системи в умовах мінливого зовнішнього середовища, зменшити вимоги до персоналу, і витрати на його навчання; підвищити автономність роботи.

РОЗДІЛ 2.

МОДЕЛІ І МЕТОДИ РОЗРОБКИ ІННОВАЦІЙНИХ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ ПРИЙНЯТТЯ РІШЕНЬ

2.1 Моделювання штучних нейронних мереж для розв'язання економічних задач

Проблеми використання штучних нейронних мереж (ШНМ) для ідентифікації економічних систем постають у зв'язку із підвищенням обчислювальних можливостей інформаційних систем економічних об'єктів і наявною потребою в ідентифікації стану останніх та охоплюють теоретичні і практичні аспекти їх застосування.

Незважаючи на великий обсяг досліджень в області ШНМ (так, С. Хайкін в узагальнюючому дослідженні «Нейронные сети. Полный курс» [232] посилається майже на 1200 найважливіших робіт з теорії нейронних мереж), багато проблем так і не отримали вичерпного рішення. На теоретичному рівні серед них виділимо наступні:

Проблема 2.1: Вибір найбільш ефективного інструменту моделювання;

Проблема 2.2: Визначення оптимальної архітектури ШНМ;

Проблема 2.3: Визначення достатнього обсягу навчальної вибірки;

Проблема 2.4: Відбір найбільш значущих вхідних даних;

Проблема 2.5: Підвищення різноманітності вибірки.

Розглянемо докладніше, що являє собою кожна із проблем (2.1-2.5).

Для проблеми 2.1 – вибору найбільш ефективного інструменту моделювання справедливе наступне твердження.

Твердження 2.1.

Ефективність вирішення складних економічних задач залежить від методів і інструментів, які застосовуються для цього.

Доведення твердження 2.1.

Відповідно визначенню 1.5 для складної економічної задачі неможливе визначення гарантовано-кращого рішення за припустимий час. Тому всі методи вирішення таких задач відшуковують тільки приблизні варіанти рішень. Отже ефективність застосування різних інструментів, зокрема нейромережевих, в загальному випадку може відрізнятись.

Твердження 2.1 доведено.

Експериментальне доведення твердження 2.1 зроблено в п. 3.1.

Сутність проблеми, яка впливає з твердження 2.1, полягає в тому, щоб визначити спосіб вибору найбільш ефективного інструменту для кожного класу економічних задач, або довести, що таке визначення в загальному випадку неможливо.

Проблема 2.2 – визначення оптимальної архітектури ШНМ впливає з відомої залежності ефективності навчання і роботи ШНМ від кількості нейронів в прихованому шарі, що особливо проявляється при нестачі вхідних

даних. В даний час ця проблема вирішується перебором різних конфігурацій і параметрів настройки ШНМ, в деяких межах, визначених переважно інтуїцією розробника. *Суть проблеми полягає в розробці формальних підходів до визначення оптимальної архітектури ШНМ, або обмеженні простору припустимих параметрів архітектури.*

Проблема 2.3 – визначення достатнього обсягу навчальної вибірки порушується в працях багатьох дослідників. Зокрема, в [32, 34, 38] доведено зв'язок між кількістю вільних параметрів ШНМ і кількістю прикладів, які потрібні для їхнього навчання, а також наближено визначено порядок останньої величини. Однак на практиці цими методами можна отримати лише досить розпливчасті оцінки, чого не завжди достатньо. *Тому суть даної проблеми полягає в розробці системи методів оцінки адекватності обсягу вхідній вибірки даних для обраної архітектури ШНМ.* Окремим випадком проблеми є аналіз розв'язуваної задачі на предмет можливості застосування ШНМ за такими ознаками, як відповідність кількості прикладів навчальної вибірки (векторів вхідної множини даних) кількості вхідних і вихідних параметрів (розмірності вхідний і вихідний вибірки).

Проблема 2.4 – відбору найбільш значущих вхідних даних безпосередньо пов'язана із зазначеною вище залежністю між кількістю вільних параметрів ШНМ і кількістю прикладів, які потрібні для її навчання. Ця проблема виникає в тих випадках, коли кількості прикладів у вхідній вибірці недостатньо для адекватного навчання нейронної мережі. *Суть проблеми полягає в знаходженні способу ранжирування вхідних параметрів за значущістю для знаходження ефективного рішення в нейромережевій моделі.*

Твердження 2.2.

Проблеми відбору значущих даних, визначення достатнього обсягу навчальної вибірки і оптимальної архітектури ШНМ повинні вирішуватися комплексно.

Доведення твердження 2.2.

Проілюструємо зв'язок між зазначеними в твердженні проблемами та ефективністю ідентифікації стану системи за допомогою графіка, наведеного на рис. 2.1 (дані умовні). Критерієм ефективності ідентифікації стану системи покладемо частку правильно розпізнаних прикладів на тестовій вибірці даних.

Крива 1 на рис. 2.1 описує зміну ефективності ідентифікації тестових прикладів при фіксованій і гарантовано достатній кількості прикладів у навчальній вибірці, в залежності від частки параметрів вхідної множини даних, які включені в цю вибірку. При цьому параметри відсортовані за зниженням значущості. Графік цієї кривої при збільшенні кількості параметрів у загальному випадку є неспадним.

Крива 2 відображає зміну ефективності ідентифікації тестових прикладів при фіксованій і гарантовано недостатній кількості прикладів у навчальній вибірці, а також за умов рівної значущості всіх її параметрів. Починаючи з певного моменту, ефективність ідентифікації при малій вибірці даних буде зменшуватися за рахунок ефекту перенавчання.

Крива 3 відображає ефективність ідентифікації тестових прикладів при

фіксованій і гарантовано недостатній кількості прикладів у навчальній вибірці, але якщо параметри відсортовані за значущістю.

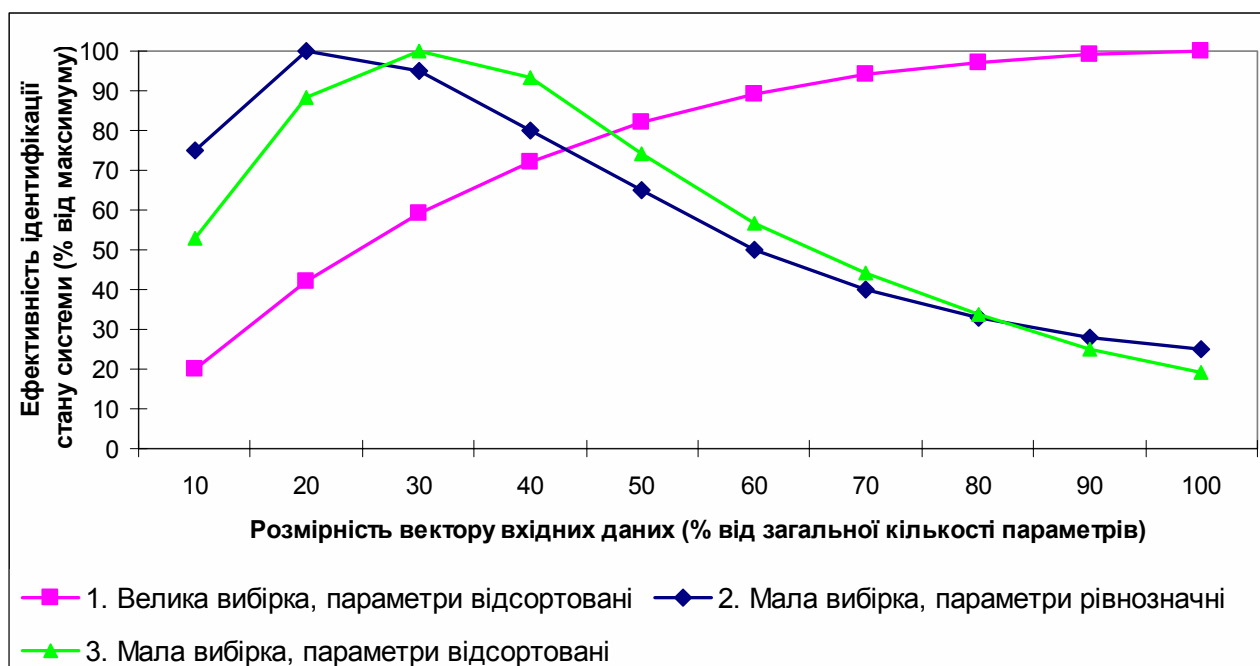


Рисунок 2.1 – Ілюстрація задачі пошуку оптимальної розмірності вектора вхідних даних

Якщо об'єктивна оцінка значущості параметрів вхідної множини даних не є можливою рекомендується обирати вибірку, розмірність вектору вхідних даних якої відповідає точці перегину графіка 2 [232]. Разом з тим як буде доведено в п. 4.2, для загальної ефективності нейромережевій моделі більшу значущість має точка перегину кривої 3 (рис. 2.1), оскільки використання більш значущих параметрів дозволяє дещо компенсувати втрату ефективності внаслідок перенавчання.

Таким чином, проблеми відбору значущих даних, визначення достатнього обсягу навчальної вибірки і оптимальної архітектури ШНМ повинні вирішуватися комплексно.

Твердження 2.2 доведено.

Проблема 2.5 – підвищення різноманітності вхідних даних виникає в тому випадку, коли навчальна вибірка складається з великої кількості векторів малої розмірності. Причому отримати прийнятне за ефективністю рішення на такій вибірці неможливо. У цьому випадку причиною може бути недостатня різноманітність вхідних даних. *Суть проблеми полягає в знаходженні методів підвищення різноманітності вхідних даних, застосування яких дозволить підвищити ефективність вирішення задачі, у порівнянні з використанням необроблених даних.*

Зрозуміло, проблеми застосування штучних нейронних мереж не вичерпується тільки проблемами 2.1-2.5. Актуальність зберігають задачі пошуку ефективних алгоритмів навчання ШНМ, інтерпретації отриманих в процесі навчання знань, підвищення швидкості навчання і багато інших,

включаючи особливості вирішення локальних економічних завдань. Однак рішення цих проблем виходить за рамки цього дослідження.

Постановка завдання і вибір ефективних інструментів його рішення.

Аналіз проблеми вибору найбільш ефективного інструменту нейромережевого моделювання доцільно починати з розгляду ролі постановки завдань для вибору інструментів їх вирішення, а також розгляду зв'язку між постановкою та ефективністю отриманих результатів. Лише в кількох джерелах [110, 143] даються рекомендації щодо застосування певних нейромережевих інструментів для вирішення певних економічних задач, наприклад, задачі побудови рейтингів, прогнозування, оптимізації.

У той же час, ще А. Лінкольн стверджував, що правильно поставлена ціль – половина успіху [91, с. 507]. Однак згодом таке розуміння значення постановки задачі збереглося тільки в соціальних науках. При вирішенні ж практичних задач під постановкою задачі найчастіше розуміють тільки формалізований опис технічного завдання. Наприклад, в «Фінансовому словнику» цей термін тлумачиться, як «точне формулювання рішення задачі на комп'ютері з описом вхідний і вихідний інформації» [227].

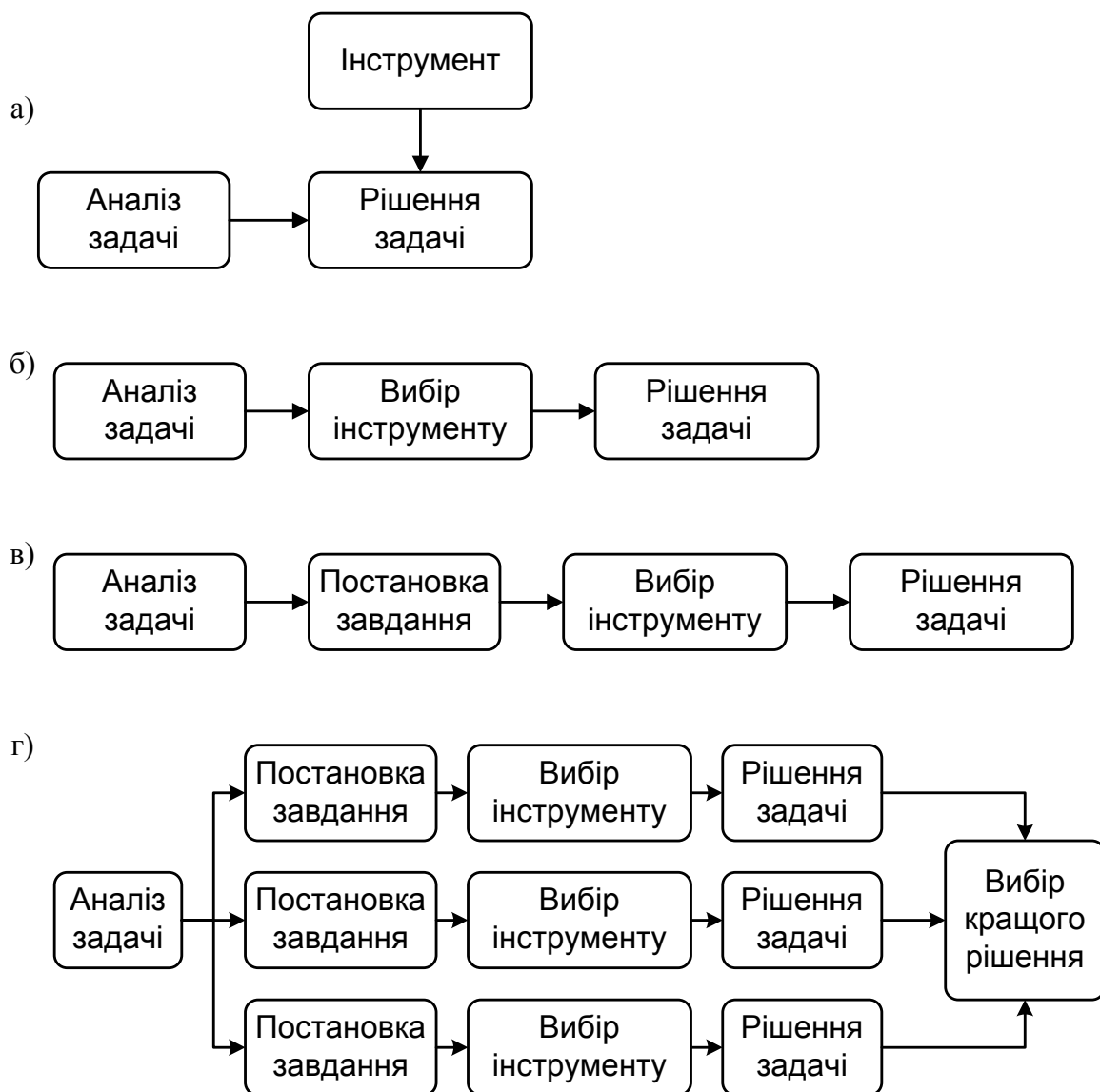
Аналіз використання цього терміна в різних літературних джерелах показує, що найбільше значення визначенню правильної постановки надається в рамках теорії розв'язання винахідницьких завдань (ТРВЗ), запропонованої і розвиненою Г. С. Альтшуллером і його послідовниками [79, 78]. Одним з інструментів, що застосовуються в ТРВЗ, є перехід від вихідної постановки задачі, що містить протиріччя, яке заважає отриманню ідеального рішення, до постановки, в якій це протиріччя усунуто. Зрозуміло, остаточна постановка також повинна забезпечувати виконання вимог замовника. Ефективність методів та інструментів ТРВЗ доведена їх використанням в даний час в більшості великих інноваційних компаній, включаючи Intel, Motorola, Ford, Philips, Siemens, Nippon, Xerox [78, с. 390], що дозволяє зробити висновок про ефективність цього підходу.

При розробці інноваційних інтелектуальних систем прийняття рішень, постановка задачі відзначена, як складова частина процесу моделювання φ (рис. 1.13). Результати проведеного дослідження показали, що при використанні інтелектуальних методів пошуку одна й та сама задача в деяких випадках може бути поставлена по-різному [161]. При цьому ефективність її рішення знаходиться в сильній залежності від постановки. Виявлення цієї залежності і формалізація підходів до постановки задач, таким чином, є актуальним завданням економіко-математичного моделювання.

На рис. 2.2 показано підходи, які використовуються, або можуть використовуватися при вирішенні економічних задач різних типів – аналізу, даних, обробки даних, чи оптимізації.

Підхід, який показано на рис. 2.2.а, є одним з найбільш поширених, не дивлячись на його простоту. У випадку 2.2.а дослідник підлаштовує задачу, що ставиться під певний інструмент її рішення. Такий підхід може бути обумовлений недостатніми навичками володіння дослідником іншими методами, або відсутністю можливостей їх реалізації в рамках наявного у

дослідника програмного забезпечення. Незалежно від причин використання такого підходу, при його використанні оптимальне рішення може бути отримано лише випадково.



а) наївний; б) інтуїтивний; в) заснований на базових постановках; г) синтетичний.

Рисунок 2.2 – Підходи до вирішення економічних завдань з використанням різних інструментів моделювання

На рис. 2.2.б, показаний підхід, який передбачає наявність в арсеналі дослідника декількох інструментів моделювання, один з яких вибирається на підставі таких критеріїв, як досвід або інтуїція.

При використанні підходу, показаного на рис. 2.2.в, до попередньої схеми додається етап постановки завдання, на якому визначається, до якого класу належить розв'язувана проблема і відповідно обирається інструмент моделювання.

Підхід, показаний на рис. 2.2.г, заснований на принципах морфологічного синтезу, описаних в 1.3. Цей підхід є найбільш трудомістким, але в той же час

дозволяє знайти найкраще рішення. Його застосування виправдане в тому випадку, коли задача може бути віднесена одночасно до різних класів і відповідним чином поставлена. Причому аналітично неможливо визначити, який з варіантів постановки призведе до найкращого рішення.

Таким чином, необхідною умовою використання підходів, поданих на схемах 2.2.в і 2.2.г, є віднесення економічної задачі до одного чи декількох з відомих класів, наприклад, відповідно до класифікації, запропонованої в 1.2. При цьому очевидно, що існує деяка відповідність між класами задач і методами їх вирішення. Розглянемо цю тезу на прикладі задач аналізу даних, класифікація яких є найбільш розгалуженою (див. рис. 1.6), і методів рішення, заснованих на використанні різних типів ШНМ.

В аналізі даних виділяються наступні класи задач: класифікації, регресії, кластеризації, аналізу зв'язків і аналізу відхилень. У задачах класифікації потрібно віднести вхідний приклад до одного із заздалегідь визначених класів. При цьому на виходах ШНМ можуть бути тільки бінарні значення, а кількість виходів відповідає кількості класів. У задачах регресії потрібно знайти залежність між вихідним параметром і набором вхідних. При цьому вихідні значення, як правило, є безперервними. Для вирішення завдань регресії використовуються багатошарові перцептрони. Вони ж можуть використовуватися для вирішення задач класифікації малої і середньої складності. Задачі класифікації високої складності (наприклад, розпізнавання образів в режимі реального часу) в даний час ефективно вирішуються за допомогою нейронних мереж глибинного навчання і згорткових ШНМ.

Сутність задачі кластеризації полягає в розбитті множини вхідних даних на кластери, що містять схожі приклади. Причому кількість кластерів з початку не задається, а сама задача вирішується за допомогою самоорганізаційних нейронних мереж, алгоритмами «навчання без учителя». Більш вузькоспеціалізованим інструментом вирішення задач кластеризації можна вважати рекурентні нейронні мережі, в яких одні і ті ж нейрони фактично є і вхідними, і вихідними, і прихованими. Такі ШНМ також ефективно вирішують асоціативні задачі в рамках класу «Аналіз зв'язків», задачі, пов'язані з очищенням даних від шуму (фільтрацією) і деякі різновиди задач оптимізації.

Пошук аномалій повинен дозволяти знаходити дані з сильними відхиленнями від основних закономірностей розподілу. Для цього можуть використовуватися самоорганізаційні ШНМ, які дозволяють наочно відобразити як основну групу даних, так і віддалені від неї об'єкти.

Найбільш ефективним інструментом виділення значимих ознак в даний час є ШНМ глибинного навчання. Основними обмеженнями їх застосування є складність реалізації і високі вимоги до обсягу навчальної вибірки.

Напрями застосування різноманітних типів ШНМ для задач аналізу даних показано на рис. 2.3.

Існують і інші підходи до класифікації задач, які вирішуються із застосуванням штучних нейронних мереж. Так С. Хайкін, відштовхуючись від базових можливостей різних типів ШНМ, виділяє задачі асоціативної пам'яті, розпізнавання образів, апроксимації функцій, управління та фільтрації [232].

Однак ця класифікація носить загальний характер і стосовно завдань економіко-математичного моделювання вимагає уточнення.

З твердження 2.1 випливає, що в загальному випадку складна економічна задача може бути вирішена різними способами. При цьому ефективність отриманих рішень може виявитися різною. Далі, на прикладі задачі біржового спекулянта, розглянемо процес її вирішення через проміжну процедуру приведення задачі до одного з розглянутих вище класів.

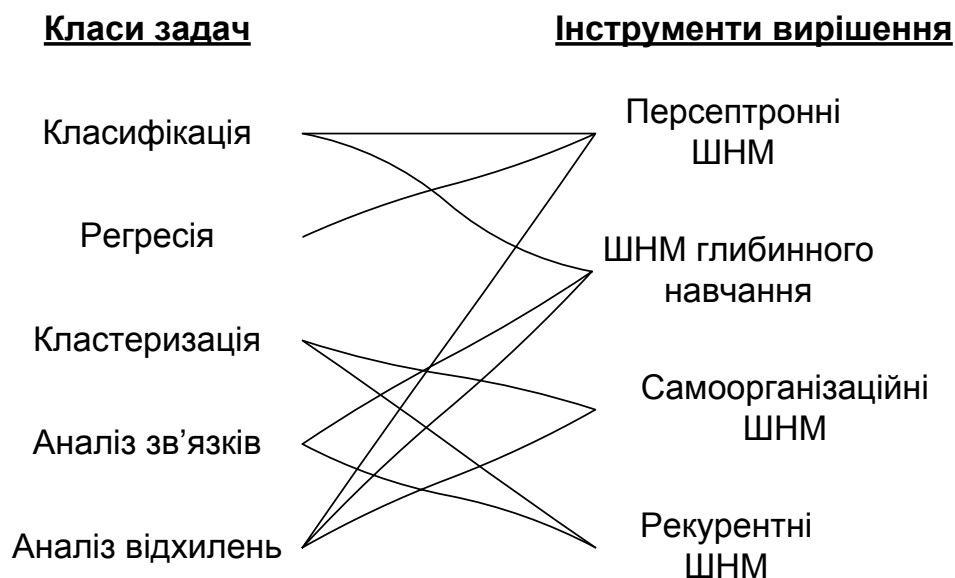


Рисунок 2.3 – Можливості застосування різних типів нейронних мереж для вирішення задач аналізу даних

У найзагальнішому вигляді задача біржового спекулянта формулюється так: необхідно розробити економіко-математичну модель, яка дозволяла б на підставі даних передісторії про зміну валютних курсів організовувати роботу на біржі (покупку і продаж фінансових інструментів) таким чином, щоб забезпечити отримання стійкого доходу.

При фіксованій стратегії роботи на біржі задача біржового спекулянта зводиться до прогнозування ринкової ситуації.

Покажемо, що задачу біржового спекулянта можна привести мінімум до трьох різних базових постановок – класифікації, регресії і кластеризації.

Базовий набір вхідних даних, які описують стан валютного ринку, в загальному випадку містить інформацію про курсові відносини різних пар валют. При цьому слід враховувати, що у всіх випадках базовий набір вхідних даних залишається без змін і містить стандартну біржову інформацію про ціни на валютні інструменти в різні часові періоди.

Розглянемо рішення цієї задачі, як *задачі класифікації*. У цьому випадку від нейронної мережі потрібно класифікувати поточну біржову ситуацію, в залежності від дій, які рекомендується зробити біржовому спекулянту – продати валюту, купити валюту, або нічого не робити. Схематично це показано на рис. 2.4.

Оскільки в даному випадку передбачено три варіанти дій, мова йде про тернарну класифікацію, яка є різновидом полінарної (див. рис. 1.6). Аналіз якості отриманих прогнозів на перевірочній вибірці дає підстави в цілому вважати отримані результати позитивними. Разом з тим, при вирішенні задачі виникають складнощі з формалізацією горизонту прогнозування. Так, при реальній роботі на валютному ринку оптимальні рішення для короткострокової, середньострокової і довгострокової торгівлі можуть бути різними. Тому при побудові моделі необхідно прийняти ряд припущень, серед яких фіксація горизонту прогнозування і підвищення порога активації вихідних нейронів. В результаті загальний економічний ефект від застосування моделі буде істотно менше максимально можливого [135].

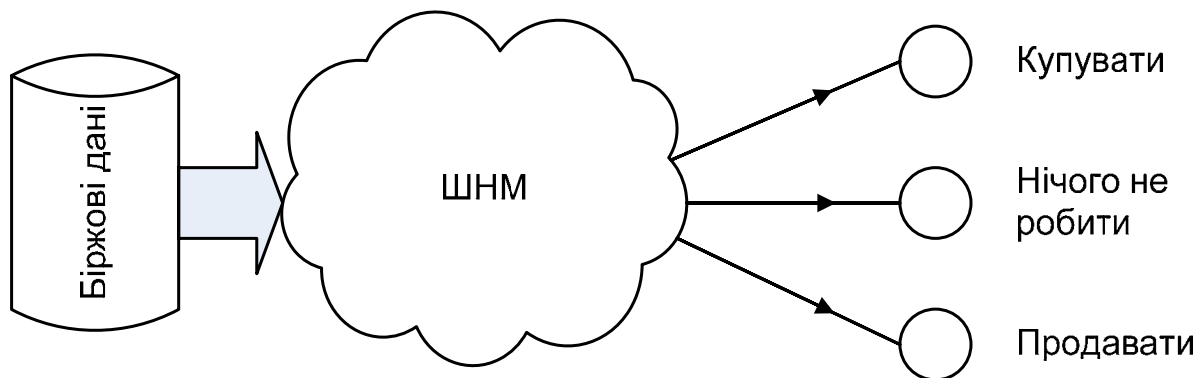


Рисунок 2.4 – Рішення задачі біржового спекулянта, як задачі класифікації

При вирішенні задачі біржового спекулянта, як задачі *регресії*, результатом роботи моделі є чисельний прогноз величини валютного курсу в наступний період, або прогноз величини відхилення валютного курсу в наступному періоді від поточного значення (рис. 2.5).

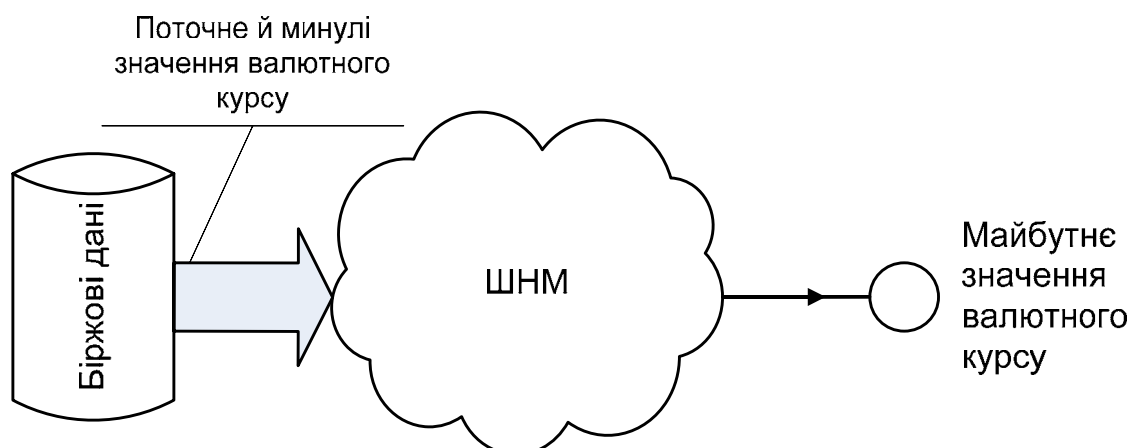


Рисунок 2.5 – Рішення задачі біржового спекулянта, як задачі регресії

Технологія вирішення цієї задачі у регресійній постановці мало відрізняється від її вирішення, як задачі класифікації. Основні відмінності виявляються в інтерпретації отриманих даних. Так, результати, що отримані в

постановці класифікації, більше підходять для створення автоматичних торгових систем, а результати, отримані в постановці регресії – для створення автоматизованих систем підтримки прийняття рішень.

Розглянемо постановку і рішення даної задачі, як *задачі кластеризації* [163]. Схематично така постановка задачі показана на рис. 2.6.

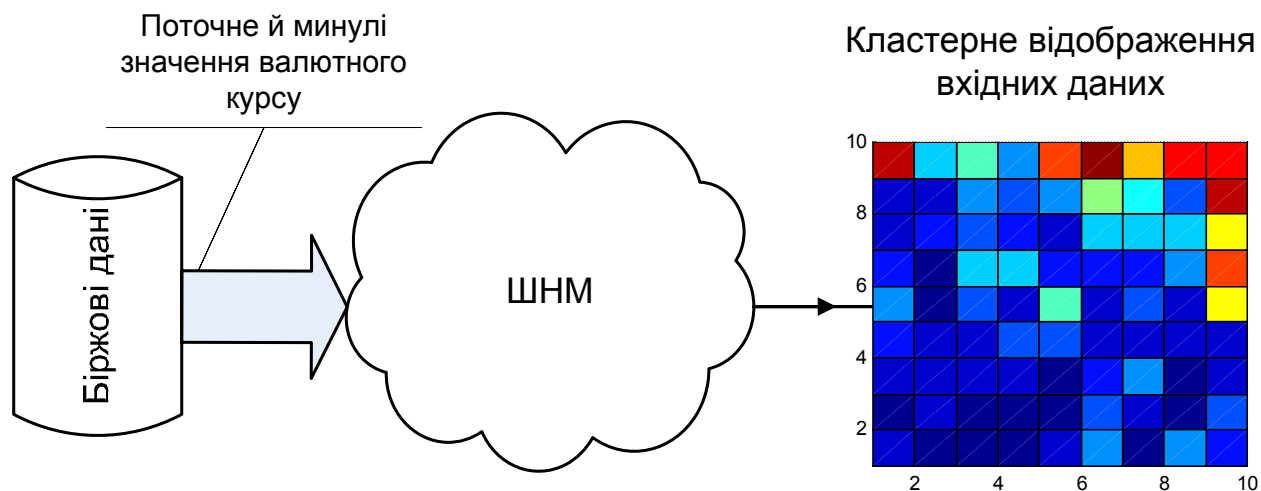


Рисунок 2.6 – Рішення задачі біржового спекулянта, як задачі кластеризації

Приводом до формулювання такої постановки є нездатність нейромережових моделей, які є результатом вирішення завдань класифікації і регресії, до аналізу ступеня детермінованості поточної ринкової ситуації. Інакше кажучи, за допомогою таких моделей можна дізнатися, в якому напрямку буде змінюватися ринковий курс, але важко визначити ймовірність такого розвитку подій. Рішення задачі біржового спекулянта, як задачі кластеризації, дозволяє розбити простір можливих наслідків на будь-яку розумну кількість кластерів і в залежності від того, до якого з них ШНМ віднесе поточну ситуацію визначати ймовірність і математичне очікування змін в ту, чи іншу сторону. Практичне розв'язання задачі, яке зроблено в п. 3.1 дозволяє отримати кращі результати, ніж при її вирішенні в постановці класифікації.

Таким чином, проведені дослідження дає підстави вважати, що ефективність отриманих результатів безпосередньо пов'язана з постановкою задачі. Це дозволяє вдосконалити методи вирішення проблеми вибору оптимальних інструментів нейромережового моделювання. Зокрема запропоновано використовувати при постановці задачі проміжний етап – зведення розв'язуваної задачі до однієї або декількох базових постановок. Остаточний вибір способу вирішення конкретної задачі може бути зроблений на підставі аналізу результатів, отриманих при її вирішенні в різних базових постановках.

Рішення задачі біржового спекулянта в запропонованих постановках, в умовах, що забезпечують порівнянність отриманих результатів розглядається в р. 3.

Методи відбору вхідних даних і архітектури нейронних мереж. Проблеми визначення оптимальної архітектури ШНМ, достатнього обсягу навчальної вибірки і відбору вхідних даних знаходяться в тісному взаємозв'язку і повинні розглядатися комплексно.

Нехай T – навчальна вибірка, яка представляє собою множину даних, що складається з векторів X_i та очікуваного відклику мережі $f(X_i)$:

$$T = \{X_i, f(X_i)\}_{i=1}^I, \quad (2.1)$$

де I – кількість прикладів у вибірці, а X_i – вектор, що має вигляд:

$$X_i = (x_{i,1}, x_{i,2}, \dots, x_{i,M}), \quad (2.2)$$

де M – кількість параметрів, що складають один приклад.

Нехай NS – структура нейронної мережі, яка в загальному вигляді може бути описана таким чином:

$$NS = \langle N, L, \psi \rangle, \quad (2.3)$$

де N – множина нейронів, що складають ШНМ;

L – множина зв'язків між нейронами;

ψ – відношення інцидентності, що ставить у відповідність кожному зв'язку з множини L два нейрона з множини N .

Для повнозв'язаних багат шарових нейронних мереж прямого поширення при описі структури можна обмежитися зазначенням кількості нейронів в кожному шарі. В цьому випадку структуру (точніше – архітектуру) мережі можна записати формулою виду

$$N^i - N^h - \dots - N^o, \quad (2.4)$$

де N^i – кількість нейронів у вхідному шарі, що збігається з M – кількістю параметрів в навчальному прикладі;

N^h – кількість нейронів в кожному з прихованих шарів;

N^o – кількість нейронів у вихідному шарі, що збігається з кількістю виходів ШНМ.

Так, наприклад, формулою 5-6-3-1 описується архітектура ШНМ, що має 5 входів, 6 нейронів в першому прихованому шарі, 3 нейрона в другому прихованому шарі і 1 вихід.

У загальному вигляді проблема відбору вхідних даних зводиться до знаходження такої вибірки $\bar{T} \subset T$, щоб при заданому параметрі обсягу вибірки

I забезпечувалося достатньо хороша апроксимація функції $f(X_i)$ і відповідно рішення досліджуваної задачі. При цьому на властивості вибірки $\bar{T}$ можуть накладатися обмеження, в залежності від умов задачі.

Існування проблеми відбору вхідних даних обумовлено залежністю між складністю структури нейронної мережі і кількістю прикладів, які необхідні для її навчання. Теоретично достатню кількість прикладів в навчальній вибірці можна визначити, відштовхуючись від концепції виміру Вапнік-Червоненкіса (*VC-виміру*), введеного в [69] і відображає «обчислювальну потужність сімейства функцій класифікації, реалізованих машинами, здатними до навчання» [232]. VC-вимір також може бути інтерпретовано як «число образів, на яких штучна нейронна мережа може бути навчена без помилок для всіх можливих бінарних маркувань функцій класифікації» [232, с. 148]. Хоча точне аналітичне визначення цього параметра для конкретної архітектури ШНМ в більшості випадків неможливо, оцінки, зроблені в [34], показують, що для найбільш поширеного типу нейронної мережі, тобто мережі прямого поширення, що складається з нейронів з сігмоїдальною активаційною функцією, VC-вимір має порядок W^2 , де W – кількість вільних параметрів ШНМ. Оскільки до вільних параметрів мережі відносяться вагові коефіцієнти зв'язків і порогові значення активації нейронів, то:

$$W=|L|+|N|. \quad (2.5)$$

При цьому, якщо архітектуру ШНМ позначити формулою виду $N_1 - N_2 - \dots - N_{lay}$, де N_1 – вхідний шар, N_{lay} – вихідний, а решта – приховані шари, то:

$$|L| = \sum_{i=1}^{lay-1} N_i \cdot N_{i+1} \quad (2.6)$$

$$|N| = \sum_{i=2}^{lay-1} N_i \quad (2.7)$$

Розглянемо вимоги до кількості прикладів у вхідній вибірці на прикладі порівняно невеликий повно нейронної мережі прямого поширення, з архітектурою 5-5-1. У такій мережі $lay = 3$, отже з (2.5)-(2.7) $|L|=25+5=30$, $|N|=6 \Rightarrow W=36$. Отже для гарантованого навчання такої мережі необхідна вибірка, в якій кількість незалежних прикладів має той же порядок, що і $36^2 = 1296$. Це означає, що для найбільш якісного налаштування параметрів даної нейронної мережі для ідентифікації стану складної системи достатньо вибірки, яка містить близько 1000 незалежних прикладів. Якщо кількість незалежних параметрів буде перевищувати VC-вимір, це вже не поліпшить ефективність ідентифікації.

Інші методи оцінки (наприклад, метод, запропонований в [4]) дозволяють задавати допустимий рівень помилки мережі, що дає можливість отримати

більш оптимістичні значення розмірів навчальної вибірки при прийнятному рівні помилки, однак вимоги до її величини все одно залишаються досить високими.

З (2.6) випливає, що в повнозв'язаних нейронних мережах кількість ступенів свободи безпосередньо залежить від кількості нейронів у вхідному і вихідному шарах, а отже від розмірності вектора вхідних даних і його представлення в ШНМ. При цьому нерідко зустрічається ситуація, коли цей параметр великий, а кількість прикладів в навчальній вибірці навпаки – мало. Так, наприклад задача передбачення банкрутств у банківській системі України пов'язана з аналізом більш ніж 30 параметрів, що характеризують активи, пасиви і фінансові результати банків, що обумовлює необхідність наявності вибірки даних, що складається з приблизно 25000 прикладів (розрахунки зроблені для нейронної мережі з архітектурою 30-5-1). У той же час вибірка, що зроблена на основі банківської статистики України має обсяг на порядок менше необхідного. Аналогічна ситуація виникає і в багатьох інших випадках, коли зібрати базу даних, обсяг якої був би достатнім для нормального навчання нейронної мережі не уявляється можливим [див., наприклад, 171, 172].

Проблема відбору вхідних даних тісно пов'язана з задачею знаходження такої структури мережі NS , яка б забезпечувала мінімум помилки при вирішенні задачі. У сукупності ці дві задачі не мають і швидше за все не можуть мати аналітичних методів рішення. Не існує також і досить надійних емпіричних методів їх вирішення, тому поширеним способом є простий перебір можливих варіантів. Однак, якщо при відносно невеликих обсягах досліджуваних даних використання переборних алгоритмів цілком допустимо, то із збільшенням обсягу вибірки і розмірності вхідного вектора даних витрати часу на навчання ШНМ істотно зростають, що робить використання переборних алгоритмів недоцільним [129].

Розглянемо такі методи зниження розмірності даних, як:

- логічний аналіз даних,
- оптимізацію представлення даних,
- кореляційний аналіз
- непрямі методи аналізу значущості даних.

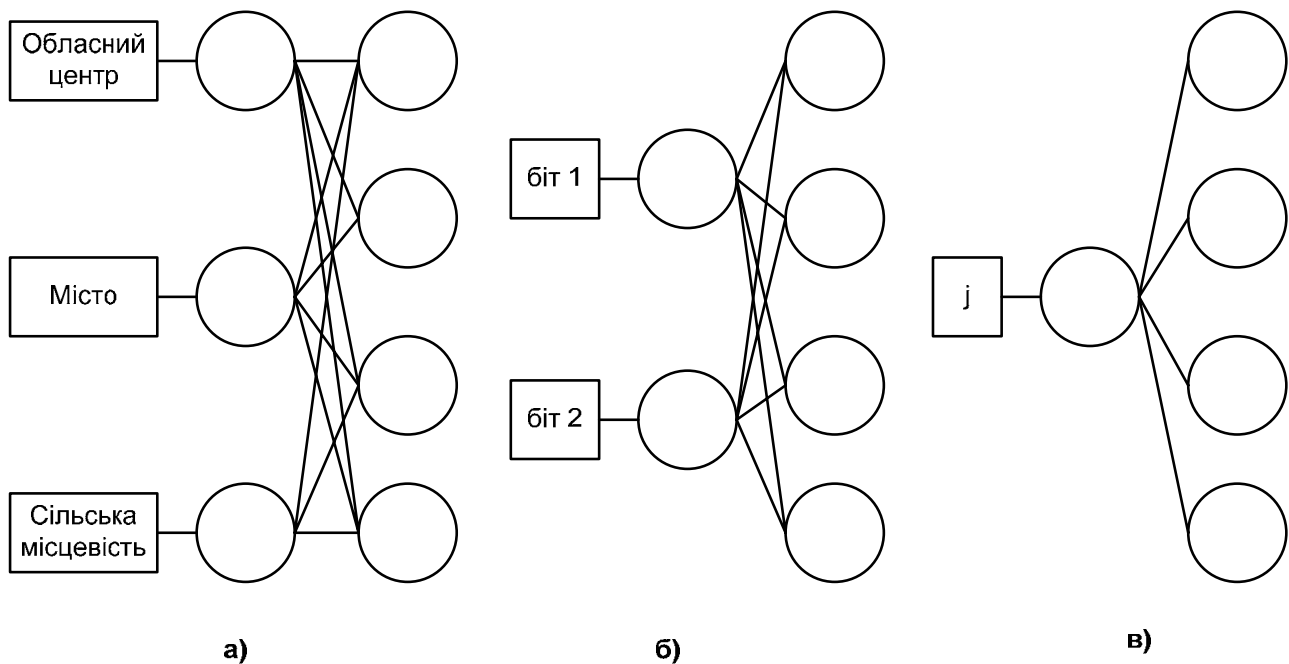
Очевидно, що основною метою застосування цих методів є позбавлення від змінних, які слабо впливають, або взагалі не впливають на вихідні параметри. Отже, в ряді випадків насамперед доцільно провести логічний аналіз даних, з метою відсікання параметрів, які не несуть інформаційного навантаження за умов задачі. Так, наприклад, при аналізі стандартних анкетних даних банківського позичальника з вибірки можна виключити поля «Порядковий номер», «№ паспорта», «ПІБ», «Адреса» і тому подібні.

Наступним методом зниження розмірності є оптимізація представлення даних. Цей метод може бути використаний для подання на вхід нейронної мережі нечислових даних. Варіанти представлення таких даних наведено на рис. 2.7.

Розглянемо переваги і недоліки кожного варіанта.

Варіант представлення даних через визначення *позиції біта* (рис. 2.7а.)

для кожного значення параметра є найбільш точним, але разом з тим і найбільш ресурсномістким, з точки зору того, скільки вільних параметрів ШНМ потрібно для його представлення.



а) представлення через позицію біта; б) представлення через бітову маску; в) представлення через ранг значення.

Рисунок 2.7 – Варіанти представлення нечислових даних, на прикладі параметра «Місце проживання клієнта»

Позначимо через v_j кількість можливих варіантів значення нечислового параметра j . При поданні параметра через позицію біта, відповідно до (2.5) - (2.7), буде потрібно

$$|J| = v_j \cdot N_2 \quad (2.8)$$

вільних параметрів, де N_2 – кількість нейронів в другому шарі ШНМ.

Для приклада, який показано на рис. 2.7а, $|J| = 3 \cdot 4 = 15$.

При кодуванні *бітовою маскою* порядковий номер кожного з варіантів представляється у вигляді двійкового числа і його біти, що мають значення «1», активують відповідні входи ШНМ (рис. 2.7б). Кількість вільних параметрів, які при цьому потрібні можна розрахувати за формулою

$$|J| = (\lceil \log_2 v_j \rceil) \cdot N_2, \quad (2.9)$$

де знак $\lceil$ позначає операцію округлення в сторону більшого числа.

Для варіанту представлення даних (рис. 2.7в), наявні значення можуть бути представлені двома бітами. При цьому значенню «Обласний центр» може відповідати маска «11», значенням «Місто» маска «10», значенням «Село» маска «01». Кількість необхідних вільних параметрів $|J| = 2 \cdot 4 = 8$.

Недоліком цього способу кодування є деяке погіршення адекватності представлення вхідних параметрів, а також складність аналізу структури ШНМ для виявлення знайдених мережею залежностей.

Представлення нечислових параметрів у вигляді *рангів значень* (рис. 2.7в) може використовуватися тоді, коли для цих параметрів існує критерій розподілу на шкалі «краще» - «гірше», але якщо такий критерій може бути синтезовано, виходячи з умов задачі. Якщо такого критерію немає, або він не є явним, даний спосіб кодування можна застосовувати тільки у виняткових випадках, що обумовлює основний недолік даного способу. Перевагою його є мінімальні вимоги до ресурсів ШНМ. Дійсно, в цьому випадку:

$$|J| = N_2, \quad (2.10)$$

що очевидно менше, ніж значення, отримані за формулами (2.8) або (2.9).

У розглянутому прикладі вхідні значення можна закодувати, наприклад, виходячи з питомої ваги проблемних кредитів в різних регіонах. Так, якщо найбільш надійними є кредити, що видані мешканцям обласних центрів, а найменш надійними – кредити мешканцям міст, то значенням «Обласний центр» відповідає код 1, значенням «Село» – код 0,33, а значенням «Місто» – код 0,66. При цьому кількість вільних параметрів складатиме всього $|J| = 4$.

Кодування нечислових даних можна розглядати як компроміс між точністю відображення вхідної інформації і використовуваними ресурсами ШНМ. При великому обсязі вхідної вибірки, або при малій кількості можливих значень параметра (наприклад, параметр «стать», який може приймати всього два значення) доцільно використовувати кодування позицією біта. При великій кількості можливих значень параметра або при невеликому обсязі навчальної вибірки доцільно використовувати кодування бітовою маскою. У виняткових випадках може бути виправдане кодування шляхом ранжирування значень. Крім того, доцільно вивчити статистичні параметри вхідних вибірки, і розглянути можливість видалення з неї рідко використовуваних значень.

Кореляційний аналіз є найбільш відомим і поширеним методом формального аналізу, на підставі якого з вхідних даних можна виключити як параметри, які занадто слабо пов'язані з вихідними даними, так і параметри, які занадто сильно пов'язані з іншими вхідними параметрами. Недоліком кореляційного аналізу є можливість виявляти залежності між зміною тільки тих параметрів, які мають числове вираження. Нечислові параметри можуть бути проаналізовані лише в тому випадку, якщо вони можуть бути розташовані за принципом «краще» - «гірше», що не завжди можливо. Сучасне програмне забезпечення дозволяє автоматизувати здійснення кореляційного аналізу, залишивши на долю дослідника підготовку даних і інтерпретацію отриманих результатів [167]. Приклад результатів автоматичного кореляційного аналізу показано на рис. 2.8.

Недоліком класичного метода визначення коефіцієнта кореляції є те, що він достовірно дозволяє знаходити тільки лінійні залежності, тоді як реальні

економічні процеси можуть розвиватися за законами, далекими від лінійних. Проведені дослідження показали, що при аналізі даних, заданих у вигляді аналітичної функції, коефіцієнт кореляції Пірсона, рівний 1 виходить тільки для лінійної залежності. Для кубічної і експоненційної залежності його значення знаходиться в межах 0.6-0.7, а для періодичних функцій, які зустрічаються в аналізі економічних систем досить часто, значення коефіцієнту кореляції близько до 0 [57].

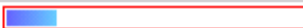
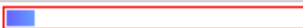
Входные поля		Корреляция с выходными полями	
№	Поле	Давать кредит	/
15	Обеспеченность займа		0,336
5	Наличие недвижимости		0,327
7	Наличие банковского счета		0,232
14	Срок проживания в данной местности, лет		0,202
9	Срок работы на данном предприятии, лет		0,201
10	Срок работы на данном направлении, лет		0,193
6	Наличие автотранспорта		0,167
13	Количество иждивенцев		0,095
12	Количество лет		0,094
8	Наличие страховки		0,080
4	Среднемесячный расход, грн		-0,073
2	Срок ссуды, мес		0,053
1	Размер ссуды, грн		0,046
3	Среднемесячный доход, грн		-0,024
11	Семейное положение		0,003

Рисунок 2.8 – Приклад результатів кореляційного аналізу анкетних даних із надійністю позичальника

Значно кращих результатів по виявленню залежностей в даних дозволяє домогтися використання сучасних методів аналізу взаємозалежностей в даних. За результатами дослідження [57] абсолютно кращим виявився метод, заснований на обчисленні коефіцієнта максимуму взаємної інформації (МІС). Значення МІС для всіх аналітично-заданих залежностей дорівнювало 1, що відповідає максимальній ступені зв'язку.

Хоча коефіцієнт МІС було запропоновано порівняно недавно, в 2011 році, він швидко набув поширення для аналізу слабкоструктурованих даних. Вже в 2012 році К. Мерфі в монографії, присвяченій перспективам розвитку машинного навчання, назвав МІС кореляцією XXI століття [47, с. 61].

Слід зазначити і недоліки МІС. Головний з них обумовлено вимогами до дискретності вхідних величин. Це змушує використовувати для аналізу безперервних величин алгоритми огрубіння даних, що негативно позначається на точності знаходження слабких залежностей. Для підвищення точності рекомендується вибір методу огрубіння проводити ітеративно, що дозволяє трохи поліпшити точність, але в свою чергу збільшує трудомісткість та витрати часу. Таким чином, використання МІС вимагає високої кваліфікації аналітика. Крім того застосування цього методу стримується його відсутністю в поширених програмних продуктах.

Важливу роль при вирішенні задачі зниження розмірності вибірки можуть зіграти непрямі методи аналізу значимості даних. В роботі [156] показано, що в якості методу відбору значущих параметрів може бути використаний будь-який метод, що дозволяє виконати ранжирування набору даних за ступенем їх впливу на вихідний параметр, навіть якщо таке ранжирування не є його основною функцією. До таких методів, наприклад, можна віднести алгоритм автоматичної побудови дерев рішень *C4.5* [55].

Для роботи алгоритму *C4.5*, обов'язковим є дотримання таких умов [55]:

- кожен елемент вхідного набору даних повинен відповідати одному з визначених класів;
- кожен приклад повинен однозначно ставитися тільки до одного з класів;
- кількість класів повинна бути значно менше кількості прикладів у вхідній вибірці.

Легко помітити, що дані умови в цілому справедливі і для будь-якої задачі, яка може бути розв'язана за допомогою перцептронних нейронних мереж.

Використання цього алгоритму дозволяє провести ранжирування не тільки тих факторів, які мають числовий вираз, але і нечислових факторів, або тих, які не можуть бути приведені до числового виду. Особливості роботи алгоритму дозволяють варіювати кількість параметрів що відкидаються (тобто таких, яким присвоюється нульова значущість), що підвищує гнучкість методу. Ще однією перевагою алгоритму *C4.5* в порівнянні з кореляційним аналізом є автоматичне відкидання параметрів, що мають сильну взаємну кореляцію з іншими.

Приклад результатів аналізу, проведеного з використанням такого алгоритму, показаний на рис 2.9.

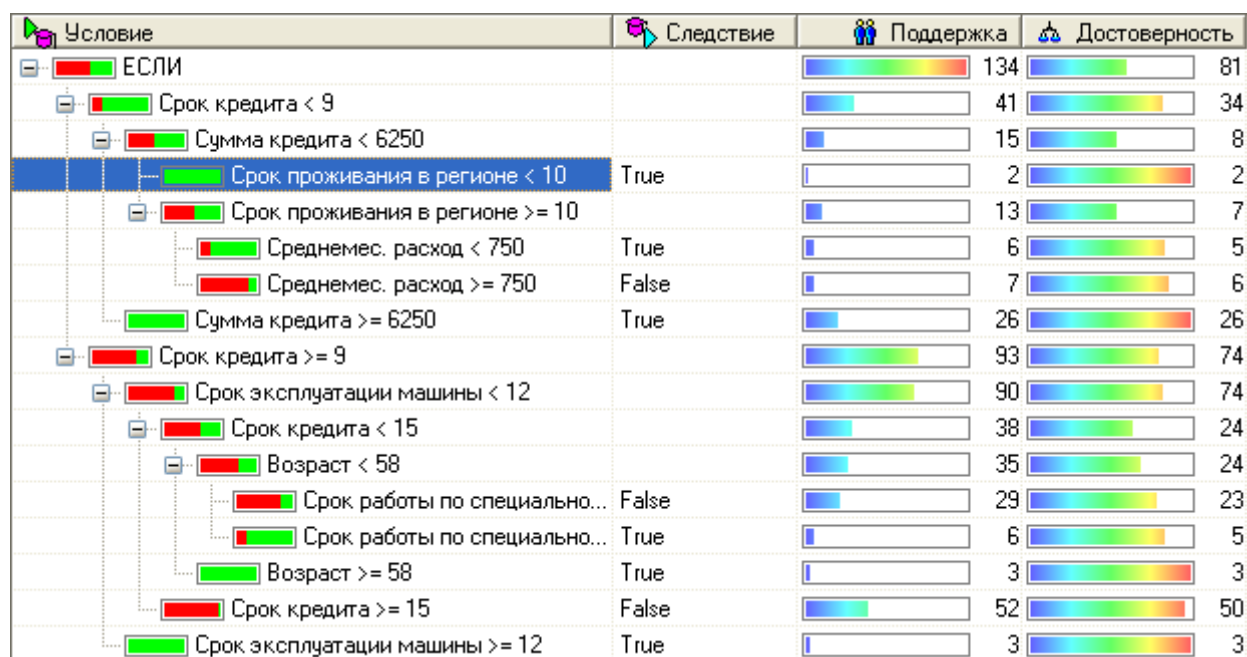


Рисунок 2.9 – Приклад результатів аналізу вибірки анкетних даних за допомогою алгоритму *C4.5*

Крім даних, які наведено на рис. 2.9, алгоритм формує також матрицю значущості вхідних параметрів. Залежно від конкретного завдання ці дані можуть бути схожими на результати кореляційного аналізу, але можуть і суттєво відрізнятися від них.

Таким чином, відбір вхідних даних є обов'язковою процедурою при невеликому обсязі навчальної вибірки. Використання непрямих методів, зокрема аналізу значущості показників за допомогою алгоритмів побудови дерев рішень дозволяє розширити інструментарій розробки нейромережових моделей в таких умовах.

Підвищення різноманітності вхідних даних при нейромережевому моделюванні. Необхідність підвищення різноманітності даних, як вже зазначалося, спостерігається при малій розмірності вектора вхідних даних, великій кількості прикладів у вибірці і наявності складної залежності між вхідними і вихідними параметрами. Прикладом таких даних може бути біржова інформація.

Базовий набір вхідних даних, що описує стан валютного ринку, в загальному випадку містить інформацію про курсові відносини різних пар валют за часовими періодами. При цьому в кожному періоді виділяється курс на початок періоду o_i , на кінець періоду c_i , а також максимальні h_i і мінімальні l_i значення курсу за період. Оскільки сам по собі такий вектор не несе ніякої інформації про тенденції змін цін на біржі, для урахування динаміки цього та інших часових рядів застосовують трансформацію вхідного набору даних за допомогою ковзного вікна, тобто замість вектора $\{o_i, h_i, l_i, c_i\}$ використовується вектор $\{o_{i-k}, h_{i-k}, l_{i-k}, c_{i-k}, o_{i-k+1}, \dots, o_{i-1}, h_{i-1}, l_{i-1}, c_{i-1}, o_i, h_i, l_i, c_i\}$. Однак і цей спосіб не дозволяє досягти високих результатів в прогнозуванні, що є наслідком з фундаментальних обмежень персептронних і подібних до них мереж, які були сформульовані Розенблаттом і доповнені М.Минським і С.Пейпертом [45]:

- персептронні мережі не здатні до узагальнення своїх характеристик на нові стимули або нові ситуації, а також не здатні аналізувати складні ситуації в зовнішньому середовищі шляхом розчленування їх на більш прості;
- персептрони мають обмеження в задачах, пов'язаних з інваріантним представленням образів.

Зменшити вплив цих обмежень можна шляхом надання персептронам додаткової інформації, що уточнює поточну ситуацію. Отже, виникає необхідність використання додаткових методів підвищення різноманітності даних.

В якості таких можуть бути використані різні емпіричні методи аналізу, які зазвичай формулюються у вигляді «умова» – «наслідок». Фактично, такі методи можна розглядати, як мікро-експертні системи, що дозволяє говорити про комбінований нейро-експертний модуль аналізу (рис. 2.10).

Розглянемо приклад побудови вхідної частини нейро-експертного модуля аналізу біржових даних.

До емпіричних методів аналізу валютних ринків зокрема відноситься прогнозування за допомогою ринкових індикаторів і осциляторів. У

найпростішому випадку осциляторні методи дозволяють отримувати вихідний сигнал у вигляді набору (-1/0/1), що відповідає прогнозуванню відповідно зниження курсу (-1), збереження нинішнього рівня (0) та підвищення курсу (1). Існують методи і для більш детального прогнозу [241].

Розглянемо принципи використання деяких типів осциляторів спільно з ШНМ. Для цього розглянемо осцилятори *ROC*, *RSI* і *MACD*. Осцилятор *ROC* (*Rate of Change*) виступає індикатором сильних рухів ринку і обчислюється таким чином:

$$ROC_t = (P_t / P_{t-n}) * 100 \%, \quad (2.11)$$

де P_t – ціна в момент часу t , n – порядок осцилятора.

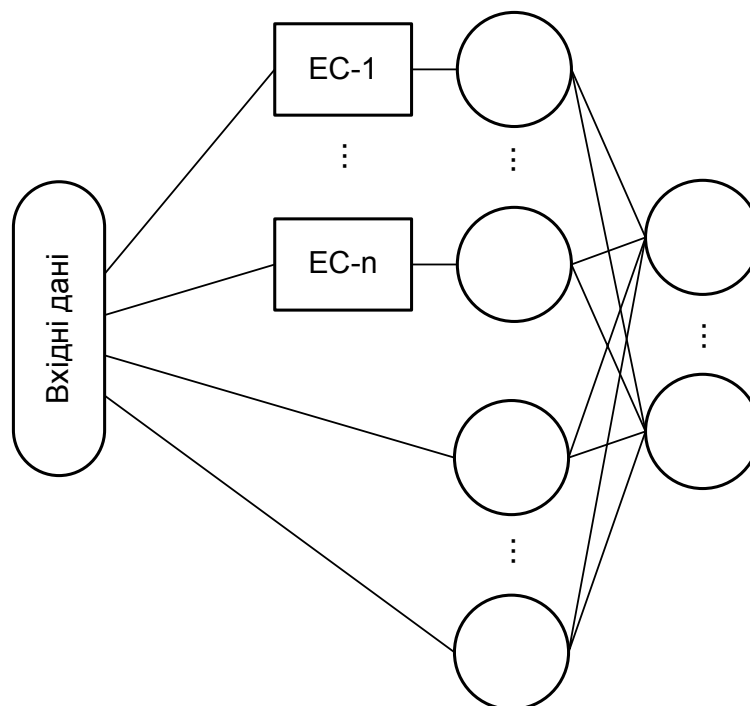


Рисунок 2.10 – Структура вхідної частини нейро-експертного модуля аналізу даних

Для прийняття рішень за допомогою осцилятора *ROC*, крім розрахованого поточного значення осцилятора, необхідно мати статистику його попередніх значень, з якої знаходиться *максимальне відхилення осцилятора*. Як правило, вважають, що осцилятор спрацював, якщо відхилення, що спостерігається, становить не менше 70-80% максимального [240]. Це дозволяє сформулювати математичну модель прийняття рішень за допомогою осцилятора *ROC* у такий спосіб:

$$D_{ROC} = \begin{cases} 1; & \left(\frac{P_t}{P_{t-n}} \right) * 100\% > H_{ROC}, \\ -1; & \left(\frac{P_t}{P_{t-n}} \right) * 100\% < L_{ROC}, \\ 0; & L_{ROC} < \left(\frac{P_t}{P_{t-n}} \right) * 100\% < H_{ROC}, \end{cases} \quad (2.12)$$

де L_{ROC} – нижня межа спрацьовування осцилятора ROC; H_{ROC} – верхня межа спрацьовування осцилятора ROC; D_{ROC} – рішення, прийняте за допомогою осцилятора ROC.

Параметри H_{ROC} і L_{ROC} задають поріг подачі осцилятором сигналів відповідно на покупку і продаж валюти. Чим ближче значення цих параметрів до максимально досягнутого відхилення, тим вірогідніше стають сигнали осцилятора, але разом з тим зростає кількість нерозпізнаних сигналів про майбутні зміни стану ринку. Тому вибір значень H_{ROC} і L_{ROC} сильно залежить від умов ринку, на якому застосовується осцилятор. На ринках з невеликими середньодобовими коливаннями цін значення цих параметрів задаються ближче до центральної позначці, на активних же ринках, з великими коливаннями цін, поріг спрацьовування навпаки збільшують.

Дія осцилятора *Relative Strength Index (RSI)* заснована на визначенні періодів знаходження ринку у стані перекупленості чи перепроданості. За гіпотезою авторів осцилятора такі періоди спостерігаються відповідно на початку спаду або підйому ціни. Модель використання осцилятора сформулюємо наступним чином:

$$Z_{RSI_t} = \begin{cases} 1; & 100 - \left(\frac{100}{1 + RS} \right) > H_{RSI}, \\ -1; & 100 - \left(\frac{100}{1 + RS} \right) < L_{RSI}, \\ 0; & L_{RSI} < \left(\frac{100}{1 + RS} \right) < H_{RSI}, \end{cases} \quad (2.13)$$

де Z_{RSI} – ознака знаходження в зоні перекупленості чи перепроданості, H_{RSI} – межа зони перекупленості, L_{RSI} – межа зони перепроданості, RS – відношення:

$$RS = \frac{AU_x}{AD_x}, \quad (2.14)$$

де x – порядок осцилятора, в днях, AU_x – кількість разів, коли ринок закрався вище цін відкриття за x днів, AD_x – кількість разів, коли ринок

закрився нижче цін відкриття за x днів [144].

Але факт ненульового значення Z_{RSI} ще не свідчить про негайну зміну тренда. Ринок може перебувати в станах перекупленості, чи перепроданості невизначений час, тому після того, як Z_{RSI} прийняло нульове значення, ОПР продовжує моніторинг ринку, але вже з метою виявлення перегину графіка RSI :

$$D_{RSI_t} = \begin{cases} 1; & RSI_t > RSI_{t-1} \cap RSI_{t-1} \leq RSI_{t-2} \cap Z_{RSI_{t-1}} = -1, \\ -1; & RSI_t < RSI_{t-1} \cap RSI_{t-1} \geq RSI_{t-2} \cap Z_{RSI_{t-1}} = 1, \\ 0; & \text{иначе,} \end{cases} \quad (2.15)$$

де D_{RSI} – рішення, прийняте за допомогою осцилятора RSI , RSI_t – значення осцилятора RSI в день t :

$$RSI_t = 100 - \left(\frac{100}{1 + RS_t} \right), \quad (2.16)$$

де t – день, для якого виконуються розрахунки.

Також осцилятор RSI може виступати у якості фільтру для проведення операцій. Так, його знаходження у зоні перепроданості забороняє операції з продажу валюти, а знаходження в зоні перекупленості – операції купівлі.

Останній з осциляторів, які розглядаються, *Moving Averages Convergence-Divergence (MACD)* засновано на різній чутливості рухомих середніх різних порядків до коливань тренду. Чутливішою при цьому виявляється рухома середня з меншим порядком. При цьому, на ринку що підвищується, більш чутлива лінія рухомої середньої розташована вище, а на ринку, що знижується, спостерігається зворотна закономірність. Легко показати, що перетин короткострокової і довгострокової рухомих середніх свідчить про зміну тенденції ринкового курсу. Математично модель прийняття рішення за допомогою осцилятора MACD запишеться так:

$$D_{MACD_t} = \begin{cases} 1; & EMA_t^S - EMA_t^L \geq 0 \cap EMA_{t-1}^S - EMA_{t-1}^L < 0, \\ -1; & EMA_t^S - EMA_t^L \leq 0 \cap EMA_{t-1}^S - EMA_{t-1}^L > 0, \\ 0; & EMA_t^S - EMA_t^L \leq 0 \cap EMA_{t-1}^S - EMA_{t-1}^L \leq 0, \\ 0; & EMA_t^S - EMA_t^L \geq 0 \cap EMA_{t-1}^S - EMA_{t-1}^L \geq 0, \end{cases} \quad (2.17)$$

де D_{MACD} – рішення, прийняте за допомогою осцилятора $MACD$, EMA_t^L – значення рухомої середньої з більшим порядком, EMA_t^S – значення рухомої середньої із меншим порядком.

Моделі (2.13), (2.15), (2.17) фактично є найпростішими експертними системами, що реалізують систему «підказок», заснованих на досвіді і інтуїції людей – експертів у зазначеній предметній області.

Крім подібних моделей для розширення вектору вхідних даних можуть використовуватися методи обробки даних, які, відповідно до розробленої в п. 1.2 класифікації, належать до групи таких, що не змінюють порядок елементів вибірки. Наприклад, для обробки біржових даних використовуються методи згладжування лінії тренда.

Використання описаних методів дозволяє істотно розширити вхідну вибірку даних. В [180] показано, що з вихідного вектора, що містить 4 показника $\{o_i, h_i, l_i, c_i\}$ може бути отриманий вектор з 20 порівняно незалежних показників, що дозволяє поліпшити якість прогнозування валютних ринків. Аналогічні способи можуть бути використані і в інших подібних випадках.

2.2 Нечітко-логічні методи вибору інструментальних засобів інтелектуальних обчислень

Моделювання інноваційних інтелектуальних систем прийняття рішень являє собою складний процес, що включає вирішення завдань різних типів і класів в різних фазах життєвого циклу проекту, серед яких виділяють [49]:

- аналіз предметної області;
- попередній збір і аналіз даних;
- збір даних;
- обробку та підготовку даних;
- моделювання;
- оцінку результатів;
- реалізацію;
- впровадження.

Для кожної з цих фаз характерні свої провідні критерії оптимальності при виборі інструментальних засобів вирішення.

Так, при аналізі предметної області і попередньому аналізі даних, основним критерієм є швидкість перевірки гіпотез. Тому інструментальні засоби повинні забезпечувати мінімальні витрати часу на побудову різноманітних моделей.

Збір і підготовка даних при вирішенні багатьох сучасних завдань пов'язані з необхідністю обробки великих масивів інформації, вимірюваних терабайтами. Отже, інструментальні засоби повинні забезпечувати можливість роботи з такими масивами.

При моделюванні важливими параметрами є точність рішення і швидкість розробки.

При реалізації та впровадженні основними критеріями стають швидкість роботи системи та потрібність в обчислювальних ресурсах. Велику роль також можуть грати фінансові витрати на необхідну кількість ліцензій програмного забезпечення і можливість інтеграції з існуючою інформаційною системою.

Крім того, важливим фактором може стати призначення ІСПР, яка синтезується. Очевидно, що набір критеріїв оптимальності для комерційної системи і системи для використання в дослідницьких, або навчальних цілях

буде відрізнятися.

Таким чином, вибір інструментальних засобів може зробити істотний вплив на процес розробки ІСПР і її ефективність. Розглянемо загальні підходи до визначення характеристик інструментальних засобів вирішення економічних задач, їх оцінки і вибору.

Для програмної реалізації ІСПР і її окремих модулів в різних фазах циклу розробки, можуть використовуватися такі основні підходи [180]:

- програмування без застосування спеціалізованих пакетів і мов;
- використання програмованих систем математичного моделювання;
- використання програмних платформ з інтерактивним інтерфейсом.

Крім того, за критерієм вартості слід виділяти [180]:

- використання вільно розповсюджуваних програмних продуктів;
- використання комерційних програмних продуктів.

Розглянемо їх докладніше.

Програмне забезпечення низькорівневої розробки передбачає безпосереднє програмування всіх алгоритмів роботи ІСПР, а також її зв'язків із зовнішніми програмами. До кінця 1990-х років це був основний спосіб розробки математичного програмного забезпечення. Тепер же такий підхід застосовується в тих випадках, коли потрібно створити програмний продукт, інтегрований в існуючу інформаційну систему. Крім того, низькорівнева розробка дозволяє забезпечити високу продуктивність. Для її досягнення можуть використовуватися розподілені обчислення, при яких основна задача розбивається на частини, які вирішуються одночасно на великій кількості ЕОМ. Недоліками такого підходу є великі витрати часу на програмування і подальше коригування програми при виникненні будь-яких змін у алгоритмі. Для розробки інформаційних систем з розвиненою математичною частиною нині знаходять застосування такі мови програмування загального призначення, як C++ та Python. Крім того існують мови, орієнтовані саме на вирішення задач аналізу даних і обробку великих масивів інформації, наприклад, мова R [141].

Значно кращу гнучкість процесу розробки забезпечує використання програмованих систем математичного моделювання. Завдяки наявності великої кількості бібліотек, що реалізують необхідні методи, при цьому значно скорочуються витрати на процес написання програмного коду. Така система надає багатий вибір математичних інструментів (який додатково може розширюватися користувачем), засобів обробки даних, засобів візуалізації. Додатково такі системи можуть надавати можливість автоматичної трансляції розроблених програм на мови загального призначення, а також створення модулів, які можуть виконуватися окремо.

Недоліки програмованих систем математичного моделювання обумовлені їхньою ідеологією, яка як вимагає від користувача певного володіння програмуванням. Найбільш відомим програмним продуктом цього класу є система Matlab [109].

Інтерактивні програмні платформи є найменш гнучким засобом розробки, оскільки надають можливість роботи тільки з певним фіксованим набором інструментів. З іншого боку, робота з системою проводиться в діалоговому

режимі і вимагає від користувача мінімальних навичок, що забезпечує високу швидкість перевірки гіпотез. Деякі інтерактивні програмні платформи (як правило, комерційні) дозволяють створювати готовий програмний продукт у вигляді модулів, здатних функціонувати окремо від системи. Як приклад інтерактивних програмних платформ можна привести Deductor Studio (бізнес-аналітика), або Neural Designer (моделювання нейронних мереж) [167].

Вибір між *комерційним* та *некомерційним* програмним забезпеченням ще порівняно недавно був простий, оскільки функціональність некомерційних програмних продуктів дозволяла використовувати їх тільки в освітніх та демонстраційних цілях. Але розвиток концепції вільного програмного забезпечення призвів до появи безкоштовних продуктів, функціональність яких впритул наближається до продуктів комерційної розробки. Прикладами можуть бути мови Scala, R і програмний каркас Apache Spark, які використовуються для обробки великих об'ємів даних. Крім того, деякі виробники комерційних програмних продуктів надають їх для освітніх цілей з мінімальними втратами функціональності (наприклад, Deductor Studio). Таким чином, в даний час є можливість використовувати якісне програмне забезпечення із мінімальними фінансовими витратами.

В узагальненому вигляді основні переваги та недоліки різних підходів до моделювання інтелектуальних систем прийняття рішень, систематизованих за рівнем спеціалізації, наведені в табл. 2.1.

Таблиця 2.1 – Аналіз використання інструментальних засобів різного рівня спеціалізації для моделювання ІСПР

Інструментарій	Переваги	Недоліки	Область використання
Мови програмування загального призначення	Висока швидкість програм; Можливість запрограмувати будь-який метод, будь-який інтерфейс і будь-який формат даних; Генерація самостійних програм	Висока трудомісткість; Необхідність знання тонкощів роботи методів і алгоритмів, мов програмування; Складність модифікації системи	Нестандартні архітектури та алгоритми навчання; Генерація самостійних програм; Інтеграція в існуючі ІС
Програмовані системи математичного моделювання	Широкий вибір готових методів, інструментів, засобів обробки даних і візуалізації; Можливість комбінувати різні методи аналізу; Генерація самостійних програм	Необхідний достатній рівень володіння програмуванням; Знання особливостей системи моделювання; Знання математичних основ методів, що використовуються	Розробка програм, що використовують інші методи аналізу; створення Комерційних програм

Інтерактивні програмні платформи	Зручний інтерфейс, широкі можливості щодо аналізу і обробки даних; Мінімальні витрати часу на навчання користувачів і процес моделювання.	Фіксований набір інструментів обробки даних і їх аналізу; обмежені можливості інтеграції існуючими інформаційними системами; низька швидкість роботи.	Розвідувальний аналіз даних; Перевірка гіпотез; Освіта; Рішення задач на невеликих та середніх обсягах даних
----------------------------------	---	---	--

Як можна побачити з даних табл. 2.1, кожен підхід має свої переваги і недоліки, значення яких може змінюватися в залежності від фаз ділового циклу. Очевидно, що при попередньому аналізі даних, для забезпечення високої швидкості перевірки гіпотез, потрібно мінімізувати витрати на процес моделювання, а при впровадженні системи в експлуатацію необхідно забезпечувати високу швидкість аналізу і обробки даних.

Важливим фактором також є необхідність в повторюваності результатів, так як критерії оптимальності для разових досліджень і для досліджень, що проводяться на постійній основі будуть істотно відрізнятися.

Методи вибору інструментальних засобів вирішення економічних завдань можуть відрізнятися ступенем участі ОПР, рівнем формалізації і трудомісткістю. Тому завдання вибору сформулюємо, як визначення найкращого варіанту з представлених інструментальних засобів, з урахуванням їх особливостей і позначених критеріїв. При цьому передбачається, що всі розглянуті засоби потенційно дозволяють вирішити поставлені завдання.

Розглянемо метод вибору інструментальних засобів, заснований на використанні павутинних діаграм.

Нехай є m варіантів вибору інструментальних засобів реалізації інтелектуальних обчислень: $p_1, p_2, \dots, p_m$. Варіанти вибору складають множину P .

Нехай є n критеріїв для порівняння інструментальних засобів в рамках вирішуваних завдань і розглянутого об'єкта дослідження: $k_1, k_2, \dots, k_n$. Набір критеріїв повинен охоплювати всі параметри інструментальних засобів, що мають значення у всіх фазах розробки ІСПР. Критерії оцінки складають множину K .

Проведемо оцінку кожного варіанту вибору інструментальних засобів за наявними критеріями. Оцінка виставляється в балах за єдиною шкалою. отримані оцінки r_{pk} запишемо в форму, яку показано в табл. 2.2.

Таблиця 2.2 – Структура таблиці для оцінювання інструментальних засобів реалізації інтелектуальних обчислень

№ _{з.п.}	Критерій	Інструментальні засоби			
		p_1	p_2	...	p_m
1	k_1	r_{11}	r_{12}	...	r_{1m}
2	k_2	r_{21}	r_{22}	...	r_{2m}
...	...	...	...	...	...
n	k_n	r_{n1}	r_{n2}	...	r_{nm}

При формуванні поля критеріїв оцінки інструментальних засобів будемо спиратися на положення стандарту ISO 9126-4:2004 «Характеристики та метрики якості програмного забезпечення» [216].

Відповідно до стандарту, якість програмних продуктів оцінюється за шістьма параметрами:

- 1) функціональним можливостям;
- 2) надійності;
- 3) практичності;
- 4) ефективності;
- 5) супроводжуємості;
- 6) мобільності.

Ці характеристики відповідно до стандарту ISO 9126-4:2004 входять в три групи - категорійно-описові, кількісні та якісні. У свою чергу в кожену характеристику входить кілька атрибутів, або субхарактеристик, які в стислому вигляді наведено в табл. 2.3.

Таблиця 2.3 – Характеристики якості програмного забезпечення відповідно до стандарту ISO 9126-4:2004

Група	Характеристика	Атрибут
Категорійно-описові метрики	Функціональні можливості	Функціональна придатність Коректність (правильність) Здатність до взаємодії Захищеність Узгодженість
Кількісні метрики	Надійність	Завершеність Стійкість до дефектів Відновлюваність Готовність
	Ефективність	Часова ефективність Ресурсомісткість
Якісні метрики	Практичність	Зрозумілість Простота використання Вивчаємість Привабливість
	Супроводжуємість	Можливість аналізу Можливість змінення Стабільність Можливість тестування
	Мобільність	Адаптованість Простота установки Співіснування (відповідність) Можливість заміщення

Всі перелічені атрибути можуть використовуватися і для оцінки інструментальних засобів реалізації інтелектуальних обчислень. Однак для того щоб уникнути надмірного ускладнення процедури оцінки, доцільно відібрати з них найбільш важливі, або провести агрегування атрибутів з близьким значенням. На підставі цих атрибутів можна сформулювати наступні критерії ефективності програмного забезпечення, відповідно до аналізу ІСПР (К.2.2.1 – К.2.2.7):

К.2.2.1. Простота використання – показує, наскільки швидко в рамках даного інструментального засобу можна реалізувати моделі з необхідною функціональністю;

К.2.2.2. Функціональна придатність – відображає наявність вбудованих математичних інтелектуальних функцій роботи з даними. Чим більше таких функцій вже реалізовано, тим краще функціональна придатність продукту;

К.2.2.3. Часова ефективність – показує, наскільки швидко в даному продукті працює реалізація інтелектуальних функцій. Для визначення швидкості доцільно зіставити час роботи різних програмних продуктів при вирішенні однієї і тієї ж типової задачі;

К.2.2.4. Ресурсомісткість – оцінка апаратного забезпечення, необхідного для реалізації поставлених завдань з використанням даних інструментальних засобів. Ресурсомісткість деяких завдань інтелектуальних обчислень вимагає їх реалізації з використанням обчислювальних кластерів, тому цей критерій має високу значимість навіть незважаючи на постійно зростаючу обчислювальну потужність сучасних ЕОМ;

К.2.2.5. Доступність використання - цей критерій агрегує атрибути «зрозумілість» та «вивчаємість». Він включає оцінку інтерфейсу програмного продукту, наявності та якості мовної локалізації. Це може мати значення при роботі з програмованими системами математичного моделювання та інтерактивними програмними платформами, більшість з яких розроблено зарубіжними виробниками програмного забезпечення.

К.2.2.6. Мобільність – агрегує всі атрибути відповідної характеристики і відображає наявність версій інструментальних засобів, сумісних з операційними системами, встановленими на комп'ютерах замовника і можливість їх інтеграції в існуючу інформаційну систему.

Крім цих атрибутів велике значення при виборі інструментальних засобів інтелектуальних обчислень має такий критерій, як:

К.2.2.7. Вартість – очікувані витрати коштів на придбання програмного забезпечення та ліцензування необхідної кількості робочих місць для реалізації поставлених завдань;

Перелічені критерії слід розглядати як основу визначення критичних параметрів інструментальних засобів інтелектуальних обчислень для конкретного об'єкта і конкретних завдань. Так, критерій наявності вбудованих математичних можливостей може бути уточнений, або розбитий на кілька додаткових, наприклад, статистичні методи аналізу, нейромережеві інструменти, засоби обробки даних і так далі.

Далі формуються підмножини критеріїв F_i , що мають значення для відповідних фаз у циклу розробки і реалізації ІСПР.

$$\forall F_i \subseteq K, \quad (2.18)$$

де i – номер фази циклу розробки і реалізації ІСПР

На заключному етапі проводиться візуальне порівняння інструментальних засобів розробки за допомогою павутинних діаграм.

Розглянемо цей метод на прикладі аналізу умовних даних.

Для множини критеріїв встановимо такі бієктивні відображення:

{Простота використання, функціональна придатність, часова ефективність, вартість, ресурсомісткість, доступність використання, мобільність} = { $k_1, k_2, k_3, k_4, k_5, k_6, k_7$ }.

Для множини фаз встановимо такі бієктивні відображення:

{Аналіз предметної області, попередній збір даних, попередній аналіз даних, підготовка даних, моделювання, оцінка результатів, реалізація, впровадження} = { $F_1, F_2, F_3, F_4, F_5, F_6, F_7, F_8$ }.

Нехай отримано наступні оцінки інструментальних засобів (табл. 2.4).

Таблиця 2.4 – Приклад оцінювання інструментальних засобів реалізації інтелектуальних обчислень

Критерії	Інструментальні засоби					
	p_1	p_2	p_3	p_4	p_5	p_6
Простота використання (k_1)	1	3	5	4	2	2
Функціональна придатність (k_2)	1	5	3	4	4	2
Часова ефективність (k_3)	5	3	2	3	5	4
Вартість (k_4)	4	1	3	2	4	5
Ресурсомісткість (k_5)	5	4	4	3	4	4
Доступність використання (k_6)	5	3	5	3	3	5
Мобільність (k_7)	5	5	5	3	4	5
Загалом балів	26	24	27	22	26	27

У порівнянні з табл. 2.2 в табл. 2.4 додано додаткове поле – «загалом балів». Воно носить інформаційний характер і в даному випадку показує, що сумарні оцінки більшості інструментальних засобів досить близькі і складають 26-27 балів.

Розглянемо процедуру вибору інструментальних засобів для двох фаз циклу створення ІСПР – попереднього аналізу даних (F_3) та реалізації (F_7).

Підмножини критеріїв вибору для них сформулюються таким чином:

$$F_3 = \{k_1, k_2, k_4, k_6\};$$

$$F_7 = \{k_2, k_3, k_4, k_5, k_7\}.$$

На підставі підмножини F_3 і даних табл. 2.3 побудуємо павутинну діаграму вибору інструментальних засобів для попереднього аналізу даних

(рис. 2.11).

Аналіз рис. 2.11 показує, що за сукупністю характеристик найбільш привабливим продуктом для попереднього аналізу даних є продукт p_3 , який лідирує за критеріями простоти і доступності використання, а також показує середні результати за критерієм функціональної придатності та вартості. Його можливими альтернативами є продукти p_4 і p_2 , у яких параметр простоти використання дещо гірше, проте забезпечується можливість доступу до більшої кількості вбудованих математичних функцій. Решту програмних продуктів використовувати недоцільно через неприпустимо низьке значення параметру простоти використання, який впливає на швидкість розробки ІСПР.

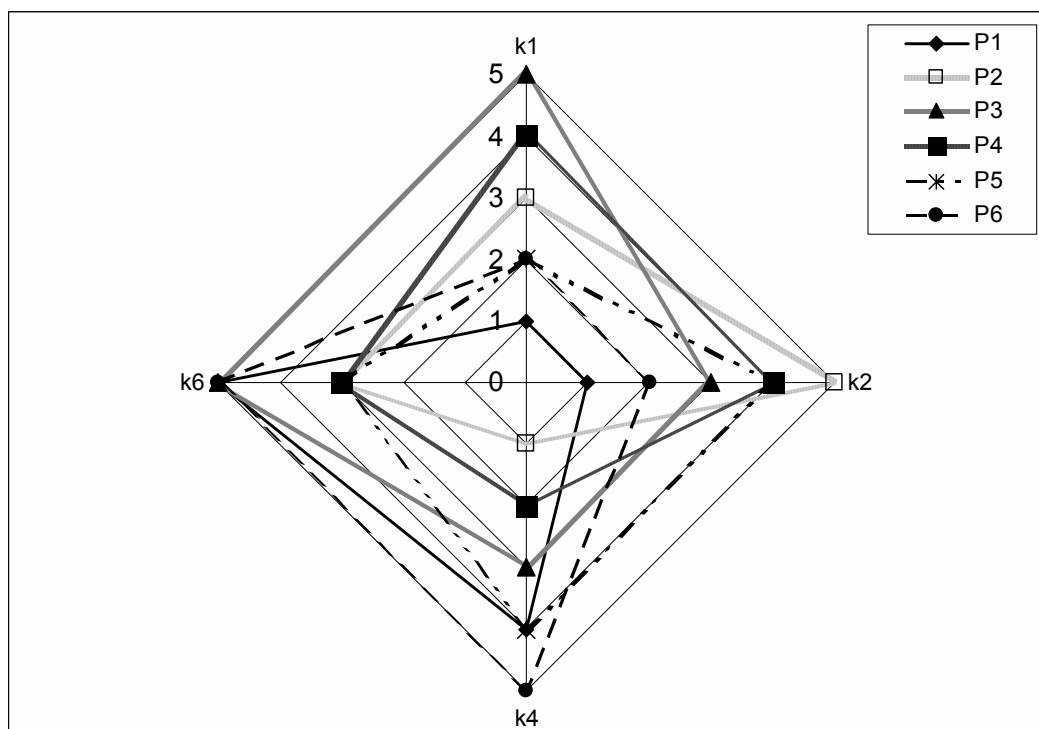


Рисунок 2.11 – Приклад діаграми вибору інструментальних засобів для попереднього аналізу даних (F_3)

На підставі підмножини F_7 і даних табл. 2.3 побудуємо павутинну діаграму вибору інструментальних засобів для реалізації ІСПР (рис. 2.12).

Аналіз діаграми на рис. 2.12 показує, що найбільш збалансованим продуктом по комплексу критеріїв є p_5 . Однак, якщо наявність великого вибору вбудованих математичних можливостей не так важлива, як критерії вартості, системних вимог, або можливості роботи в різних операційних системах, то перевагу може бути надано програмним продуктам p_1 або p_6 .

При вирішенні практичних завдань може знадобитися модифікація деяких критеріїв, або додавання нових. Так, замість загального критерію «функціональна придатність» можуть бути вказані більш чіткі вимоги, наприклад: «аналіз даних великої розмірності», «вбудовані нейромережеві інструменти», «вбудовані методи оптимізації» і так далі.

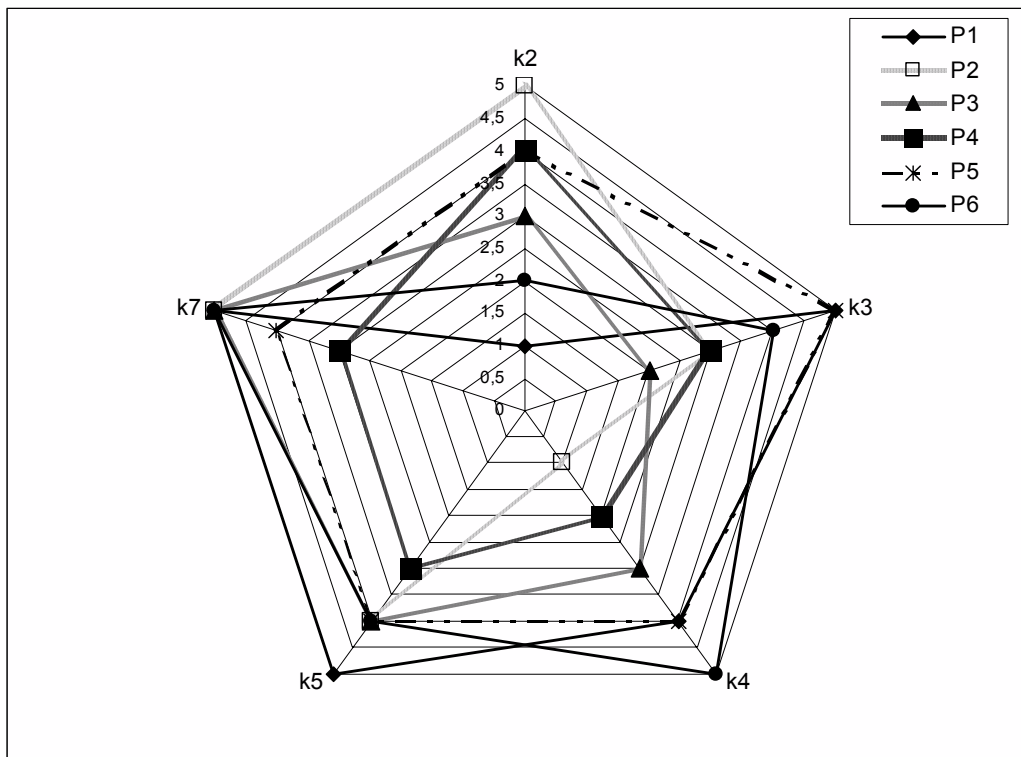


Рисунок 2.12 – Приклад діаграми вибору інструментальних засобів для реалізації ІСПР (F_7)

Метод вибору інструментальних засобів вирішення економічних задач, заснований на використанні павутинних діаграм, є достатньо наочним і простим у використанні. Проте, йому властиві деякі недоліки, серед яких слід відзначити:

- неможливість урахування різної значущості критеріїв;
- погіршення наочності при збільшенні кількості порівнюваних програмних продуктів;
- складність чисельної оцінки деяких критеріїв.

Зменшити вплив цих недоліків і вдосконалити процес вибору інструментальних засобів вирішення економічних задач можна, перейшовши до використання для формування оцінок нечітких множин. Ідея застосування нечітких множин для оцінки програмних засобів не є новою. Так, наприклад, В. Рогозін використовував нечіткі множини для оцінки привабливості web-браузерів [204]. Однак, оскільки основними характеристиками нечіткої системи є структура оцінювання, а також лінгвістичні змінні та їх параметри, використовувати результати, які отримано іншими авторами не уявляється можливим.

Рішення задачі в нечітких термінах передбачає додавання до основної процедури двох додаткових етапів – фазифікації і дефазифікації.

Відповідно базовим поняттям нечіткої логіки, нечітка множина $\bar{A}$ елементів на X визначається, як [114]:

$$\bar{A} = \{(x, \mu_A(x)) | x \in X\}; \mu_A(x) : x \rightarrow [0,1], \quad (2.19)$$

де $\mu_A(x)$ – функція приналежності нечіткої множини.

Функція приналежності може здаватися як дискретною, так і безперервною. В процесі фазифікації відбувається приведення вхідних даних, виражених в лінгвістичній формі, до термів нечіткої логіки, тобто до кожної лінгвістичної змінної ставиться у відповідність функція, яка визначає ступінь приналежності значень x цієї змінної.

У розглянутій задачі для всіх вхідних змінних нечіткої моделі використовуються кусково-безперервні функції приналежності, що складаються з п'яти термів і мають форму трапеції (табл. 2.5).

Таблиця 2.5 – Параметри приналежності критеріїв оцінки привабливості інструментальних засобів у вигляді лінгвістичних змінних

Інтервал значень	Лінгвістична оцінка привабливості продукту	Коротке позначення
[0; 0; 0.2; 0.25]	Дуже низька	ДН
[0.15; 0.25; 0.35; 0.4]	Низька	Н
[0.3; 0.4; 0.55; 0.6]	Середня	С
[0.45; 0.6; 0.7; 0.75]	Висока	В
[0.65; 0.75; 1; 1]	Дуже висока	ДВ

Аналітично трапецевидна функція приналежності, що задана на інтервалі $[a; b_1; b_2; c]$, визначається наступним чином (2.20):

$$\mu_A(x) = \begin{cases} 0, & x < a, x > c; \\ \mu_A(b), & b_1 < x < b_2; \\ \mu_A(x) \cdot \frac{x-a}{b_1-a}, & a \leq x \leq b_1; \\ \mu_A(x) \cdot \frac{c-x}{c-b_2}, & b_2 \leq x \leq c. \end{cases} \quad (2.20)$$

У графічному вигляді функція приналежності, яку визначено в табл. 2.5, показана на рис. 2.13.

Процедура оцінювання інструментальних засобів реалізації інтелектуальних обчислень частково відповідає розглянутим вище, з поправкою на використання нечіткої логіки. Так, на початку процедури проводиться оцінка кожного варіанту, на підставі чого складається таблиця, аналогічна табл. 2.2 за таким же набором критеріїв. Відмінність в тому, що оцінки програмних продуктів виставляються не в балах, а в лінгвістичних термах.

Далі, також в лінгвістичних термах вказується значимість кожного критерію оцінки для відповідних фаз. Функція приналежності задається

відповідно табл. 2.6.

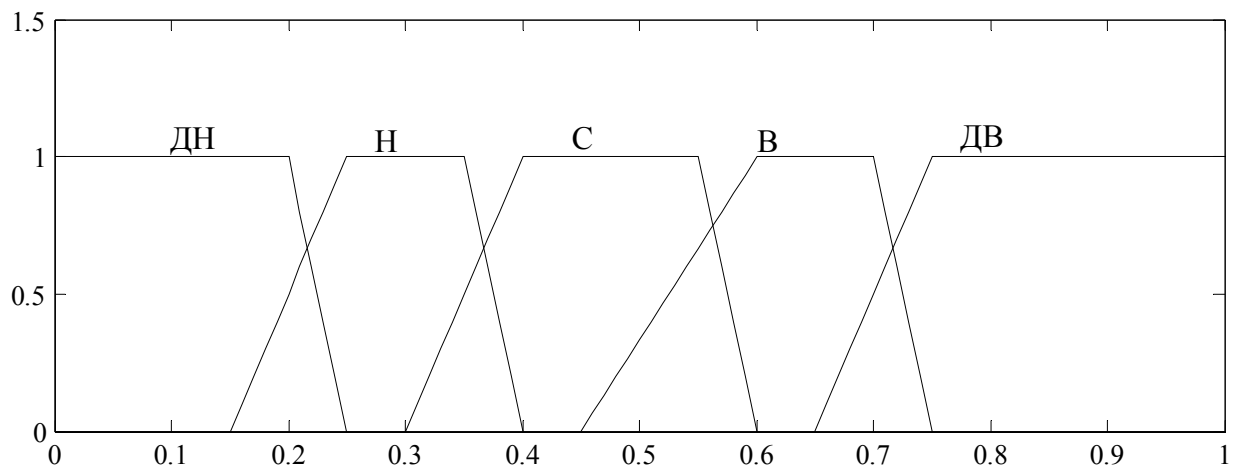


Рисунок 2.13 – Функція приналежності критеріїв оцінки

Таблиця 2.6 – Параметри приналежності значущості критеріїв оцінки у вигляді лінгвістичних змінних

Інтервал значень	Лінгвістична оцінка значущості критерію	Коротке позначення
[0; 0; 0; 0.4]	Не має значення	НЗ
[0.2; 0.35; 0.5; 0.6]	Середня	С
[0.45; 0.6; 0.7; 0.8]	Висока	В
[0.65; 0.75; 1; 1]	Дуже висока	ДВ

У графічному вигляді цю функцію приналежності показано на рис. 2.14.

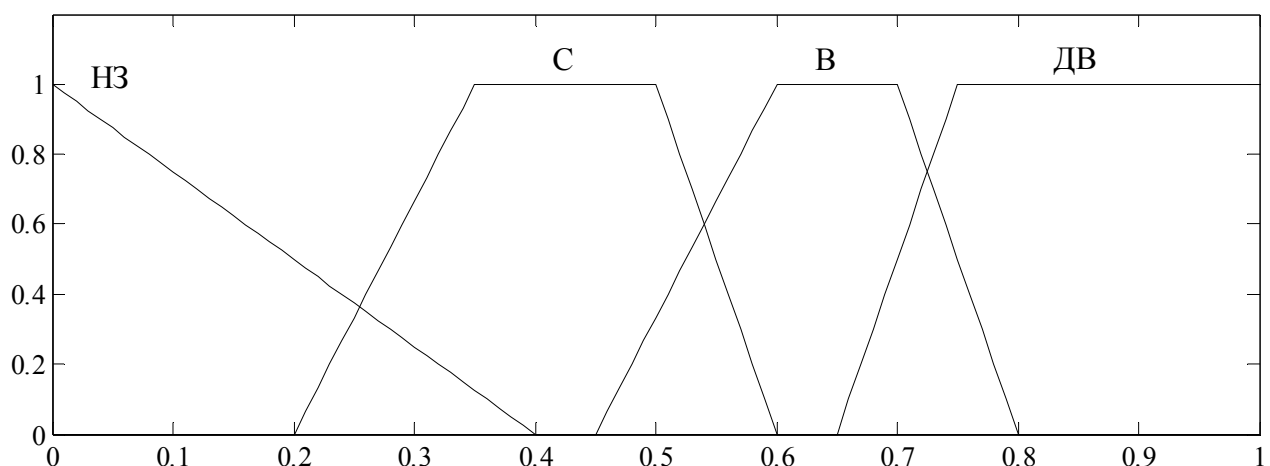


Рисунок 2.14 – Функція приналежності значущості критеріїв оцінки.

Завдяки такому вигляду функцій приналежності, використання нечітких множин може привести до результатів набагато ефективніше, ніж при використанні звичайних методів багатокритеріальних оцінок, розглянутих,

наприклад в [204]. Так, якщо два програмних продукти отримали однакові оцінки за критеріями, які мають середню, високу і дуже високу значимість, але відрізняються за критеріями, які позначені, як «не мають значення», то в традиційній системі оцінювання ці програмні продукти отримають однаковий рейтинг, а в нечіткої один з них буде оцінений вище. З погляду людини оцінка нечіткої системи є більш достовірною, оскільки як би мало не значили фактори переваги одного програмного продукту над іншим, але за інших рівних умов слід вибрати той з них, який краще за цими додатковими параметрами.

Так, як над нечіткими числами можуть проводитися операції, аналогічні традиційним арифметичним операціям, то, зокрема, результат складання нечітких чисел A і B визначається виразом [114]:

$$\begin{aligned} (A + B) &= \int_{\bar{X}} \frac{\min[\mu_A(x), \mu_B(y)]}{(x + y)}, \\ \mu_{(A+B)}(z) &= \sup_{z=(x+y)} \min[\mu_A(x), \mu_B(y)]. \end{aligned} \quad (2.21)$$

Для визначення результату множення нечітких чисел A і B служить вираз [114]:

$$\begin{aligned} (A \otimes B) &= \int_{\bar{X}} \frac{\min[\mu_A(x), \mu_B(y)]}{(x \cdot y)}, \\ \mu_{(A \otimes B)}(z) &= \sup_{z=(x \cdot y)} \min[\mu_A(x), \mu_B(y)] = \sup_{z=(x \cdot y)} \min[\mu_A(x), \mu_B(z/x)], x \neq 0. \end{aligned} \quad (2.22)$$

Використовуючи операції складання (2.21) і множення (2.22) нечітких чисел, можна визначити процедуру нечіткої оцінки інструментальних засобів вирішення економічних завдань.

Нехай P_i^j – нечітка оцінка i -го інструментального засобу по j -му критерію.

Нехай K_j^F – нечітка оцінка значущості j -го критерію для фази F циклу створення ІСПР.

Тоді нечітка оцінка RP_i^F привабливості i -го інструментального засобу для фази F може бути знайдена в такий спосіб:

$$RP_i^F = \sum_j P_i^j \otimes K_j^F. \quad (2.23)$$

Оцінки, які отримано з (2.23) можна проаналізувати шляхом порівняння їх функцій приналежності, поданих у графічному вигляді (рис. 2.15).

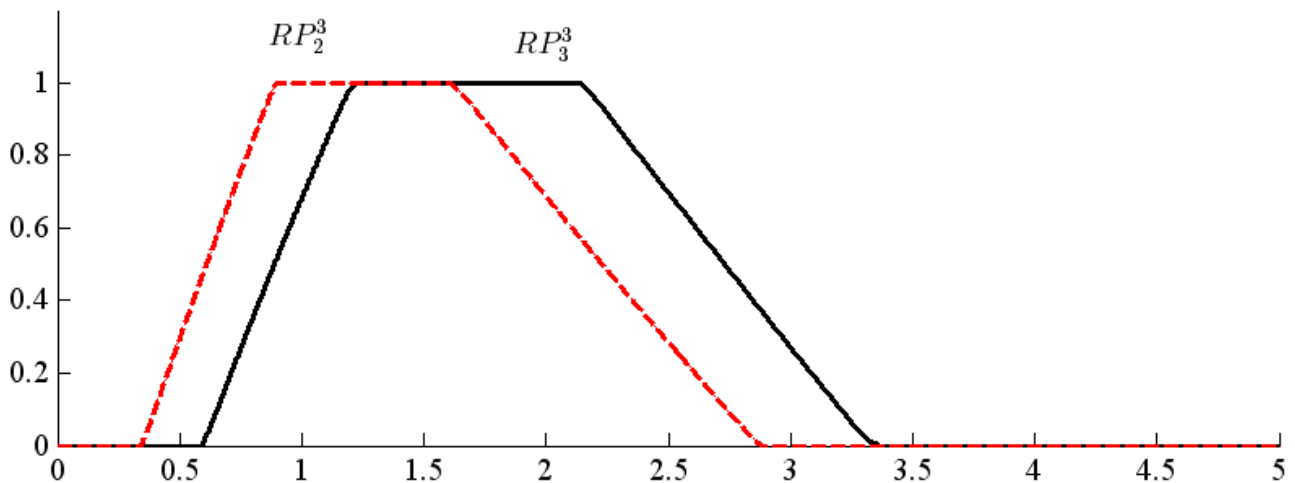


Рисунок 2.15 – Приклад порівняння функції приналежності результату оцінки

Однак більш зручним є порівняння, яке засноване на числових коефіцієнтах оцінок. Для їх отримання проводять процедуру дефазифікації.

Дефазифікація передбачає визначення правил перекладу результатів рішення в чіткі величини, які можуть використовуватися в подальших процедурах. Існує досить велика кількість різних методів дефазифікації, однак одним з найбільш поширених є метод, заснований на обчисленні центру ваги нечіткої множини. Значення координати $\bar{x}$ центру ваги визначається наступним чином [114]:

$$\bar{x} = \frac{\int_{x \in X} x \cdot \bar{\mu}_{A_i}(x) \cdot dx}{\int_{x \in X} \bar{\mu}_{A_i}(x) \cdot dx}. \quad (2.24)$$

Наприклад, стосовно нечітким множинам оцінок, функції приналежності яких показані на рис. 2.15, обчислення центрів ваги за формулою (3.24) дає значення:

$$\begin{aligned} \bar{r}p_2^3 &= 1.4570, \\ \bar{r}p_3^3 &= 1.8369, \end{aligned}$$

що можна трактувати як перевагу програмного продукту p_3 над p_2 для застосування в фазі 3 циклу створення ІСПР.

Слід зазначити, що розглянуті вище методи вибору інструментальних засобів реалізації інтелектуальних обчислень не виключають, а доповнюють один одного. Так, павутинні діаграми дозволяють швидко зіставити кілька інструментальних засобів, наочно показуючи співвідношення їх переваг і недоліків. Нечітко-логічна модель може ефективно працювати при великій кількості аналізованих продуктів і критеріїв, а також дозволяє врахувати всі параметри досліджуваних продуктів. Таким чином, павутинні діаграми можуть

використовуватися для остаточного прийняття рішень по вибору інструментальних засобів, попередньо відібраних за допомогою нечіткої логічної моделі.

2.3 Методи оцінки ефективності інтелектуальних методів вирішення економічних задач

Питання дослідження ефективності рішень природним чином виникає при розгляді економічних задач. Очевидно, що від застосування в системах прийняття рішень інноваційних методів інтелектуальних обчислень очікується отримання економічного ефекту. У той же час особливості технологій інтелектуальних обчислень зумовлюють високу варіативність отриманих результатів, що не дає можливості заздалегідь оцінити вигоди, які дає їх застосування.

Перелічені обставини зумовлюють актуальність досліджень в області оцінки ефективності інтелектуальних методів вирішення економічних задач. Складність процесів інтелектуальних обчислень викликає необхідність застосування комплексного підходу до забезпечення їх ефективності.

У широкому розумінні завдання оцінки ефективності інтелектуальних методів вирішення економічних задач може розглядатися на декількох рівнях.

Перший рівень є концептуальним. Він відповідає етапу, який пов'язаний із розробкою структури ІСПР. При цьому необхідно визначити, до якого класу належить задача, і які методи інтелектуальних обчислень можуть використовуватися для вирішення задач даного класу. Необхідно також враховувати, що одна і та ж практична задача може бути вирішена з різних позицій і відповідно буди віднесена до різних класів. Відзначимо, що на цьому рівні оцінка ефективності здійснюється до отримання фактичних результатів з практичного використання ІСПР, тобто *a priori*. Питання вибору методів інтелектуальних обчислень на цьому рівні розглянуто в 2.1.

Другий рівень відповідає оцінці різних варіантів реалізації структури ІСПР. На цьому рівні передбачається, що дослідником вже розроблено кілька моделей предметної області і необхідно вибрати кращу з них. Рішення задач оцінки, що виникають на цьому рівні, розглянемо нижче.

Н. Паклін і В.Орешков звертають увагу на два основні питання, що пов'язані з оцінкою ефективності моделей бізнес-аналітики [195, с. 563]:

- 1) Наскільки розроблена модель корисна, ефективна і точна?
- 2) Чи можна побудувати модель, яка буде працювати краще наявної?

Спроба відповісти на ці питання призводить до усвідомлення необхідності як мінімум двох підходів до оцінки ефективності інтелектуальних обчислень – абсолютного і відносного.

Визначення *абсолютних оцінок* має давати можливість проаналізувати ефективність деякої моделі і відповісти на питання про її придатність, або непридатність для вирішення поставленої задачі.

Визначення *відносних оцінок* має давати можливість порівняння

декількох моделей і визначення кращої з них для діючих умов. При цьому самі моделі можуть мати різну природу, та оперуватимуть різними наборами вхідних і вихідних даних.

Крім цього, задачу оцінки ефективності слід розглядати по відношенню до різних об'єктів і процесів, серед яких відзначимо:

- алгоритми і програмне забезпечення;
- процес навчання;
- результати аналізу даних
- результати обробки даних.

Розглянемо зміст поняття ефективності по відношенню до цих об'єктів.

Необхідність оцінки ефективності роботи програмного забезпечення обумовлено тим, що різні реалізації принципів інтелектуальних обчислень можуть сильно відрізнятися як по точності, так і за швидкістю роботи [148]. Аналогічна ситуація виникає при зіставленні ефективності роботи різних алгоритмів, що виконують функцію навчання, або налаштування параметрів інтелектуальних обчислень [63]. Таким чином, необхідний надійний метод зіставлення різних програмних продуктів, алгоритмів і ефективності їх реалізації.

Оцінку ефективності машинного навчання необхідно розглядати в зв'язку з вирішенням задачі оптимальної настройки алгоритмів навчання. В 2.1 зазначалося, що ці алгоритми чутливі не тільки до обсягу навчальної вибірки, але і до параметрів процесу навчання. Так, нейронні мережі після певного моменту, відповідного найкращої апроксимації справжніх взаємозв'язків в даних, далі починають просто запам'ятовувати вхідну вибірку даних, що погіршує ефективність їх роботи на реальних даних і носить назву «ефект перенавчання». Для дерев рішень схожа проблема виникає при збільшенні кількості розгалужень, що зменшує формальні показники помилки, але погіршує практичну значущість результатів. Існуючі підходи до оптимізації процесів навчання носять емпіричний характер і потребують узагальнення.

Оцінка ефективності результатів аналізу даних ускладнюється тим, що кожна задача в загальному випадку може мати своє трактування поняття оптимальності, яке необхідно враховувати при розробці підходів до оцінки результатів її рішення. Наприклад, в одному випадку більш важливим може статись виявлення «сприятливих» подій, а в іншому – виявлення «несприятливих». Також слід враховувати, що, принаймні, для різних класів економічних задач підходи до оцінки ефективності можуть відрізнятися. Так, методи оцінки результатів класифікації не завжди підходять для оцінки ефективності вирішення завдань регресії і навпаки. Отже, актуальним завданням є систематизація підходів до оцінки ефективності результатів інтелектуальних обчислень.

При оцінці ефективності результатів обробки даних необхідно знайти відповідь на питання про економічну вигоду, що отримується в результаті застосування різних методів обробки. Якщо економічна задача, що розглядається, безпосередньо зводиться до однієї з задач обробки даних, то вигода визначається безпосередньо. В іншому випадку необхідно

використовувати непрямі методи оцінки ефективності.

Таким чином, необхідність в оцінці ефективності виникає на таких етапах реалізації інтелектуальних методів вирішення економічних задач, як вибір програмного забезпечення, навчання інтелектуальних систем і оцінка результатів їх роботи для завдань аналізу і обробки даних. Розглянемо їх докладніше.

Оцінка ефективності роботи програмного забезпечення та алгоритмів навчання. Розвиток теорії інтелектуальних обчислень призвів до появи великої кількості модифікацій методів і алгоритмів машинного навчання. Як правило, жодна з цих модифікацій не дозволяє домогтися поліпшення ефективності у всьому діапазоні значущих параметрів і для всіх різновидів завдань, але для окремих видів приріст ефективності може бути суттєвим (див., наприклад, [63]). Аналогічна ситуація спостерігається і по відношенню до програмного забезпечення, виробники якого часто використовують спрощені реалізації методів та алгоритмів, які не дозволяють отримати кращі рішення.

Аналіз існуючих модифікацій деяких з видів машинного навчання, розглянутих вище, показує наступне.

Для побудови *дерев прийняття рішень* розроблено близько 10 різних алгоритмів, з яких актуальними, тобто такими, які знаходять застосування в даний час, є 6 (див. п. 1.2). Це означає, що для кожного з цих 6 алгоритмів існує така постановка задачі, для якої його вибір є оптимальним.

Для *генетичних алгоритмів*, в порівнянні з оригінальним варіантом, який було розроблено Холландом та Ді-Янгом [26, 15]) з'явилася значна кількість доповнень, що модифікують базові генетичні оператори. Так, тільки в системі Matlab є 5 варіантів оператора мутації, 5 варіантів оператора селекції, 4 варіанти масштабування пристосованості, 3 варіанти генерування вихідної популяції (справедливо для версії 7.7).

Найбільша різноманітність варіантів реалізації спостерігається для нейронних мереж. Вище вже зазначалося, що загальна кількість типів ШНМ в даний час перевищує 20. При цьому, наприклад, в Matlab тільки для одного з цих типів - персептрона - передбачено більше 10 варіантів навчання, серед яких кілька різновидів класичного *backpropagation*, спряжені градієнтні алгоритми, квазіньютонівські методи, алгоритм Левенберга – Марквардта.

Очевидно, що здійснення повного перебору всіх доступних варіантів для кожної задачі істотно збільшує трудомісткість реалізації ІСПР. Тому на практиці до цього вдаються рідко, покладаючись на досвід і інтуїцію розробника. Однак такий підхід не завжди дозволяє одержати кращий результат. Більш ефективним є така організація процесу вибору, при якій безліч варіантів зводиться до мінімуму за допомогою апріорного порівняння.

Питання порівняння ефективності алгоритмів і їх реалізацій вперше було піднято Д. Кнутом [124]. При цьому сам Д. Кнут посилається на ідеї А.Маркова з теорії алгоритмів, розвинені згодом Н. Нагорним [140]. Незважаючи на загальновизнану важливість цих досліджень, слід зазначити, що для оцінки ефективності роботи програмного забезпечення та алгоритмів інтелектуальних обчислень запропоновані в них ідеї повинні бути доповнені і

розвинені з урахуванням останніх досліджень.

Зокрема, основним критерієм ефективності алгоритмів в [124] приймається швидкість роботи. Однак, з огляду на те, що задачі, які відповідають рівню 2 і, особливо, рівню 3 піраміди Давенпорта (рис. 1.3) не завжди можуть сходитися до єдиного вірного рішення, іншим важливим критерієм слід вважати точність одержуваних результатів. Оскільки критерії швидкості і точності є взаємно суперечливими, в залежності від умов задачі один з них слід обмежити. Тобто можна шукати або найбільш точний алгоритм, при заданих обмеженнях за часом роботи, або найбільш швидкий алгоритм при заданих обмеженнях по точності.

Для отримання порівнянних результатів необхідно забезпечити однакові умови тестування для різних методів. З технічного боку це повинно проявлятися в однаковій апаратній платформі для перевірки різних алгоритмів. Не меншу важливість має забезпечення подібності задач для тестування. Для виконання останньої умови можна запропонувати підхід до реалізації та зіставлення алгоритмів машинного навчання, заснований на понятті *типових задач*.

Визначення 2.1.

Типовою назвемо задачу, основні характеристики якої є досить близькими для деякого набору інших задач у схожій постановці.

До задач, які можна використовувати в якості типових з економіки, будемо висувати такі вимоги (В.2.1 – В. 2.6):

В.2.1. Доступність вхідних даних, причому в різних варіантах, які, тим не менш, мали б схожі характеристики розподілу. Виконання цієї вимоги необхідно для забезпечення порівнянності результатів досліджень, проведених у різні часи і в різних умовах.

В.2.2. Прозорість економічної інтерпретації результатів. Необхідно для забезпечення можливості прямого зіставлення алгоритмів, які досліджуються, за основним критерієм, прийнятим в економіці – вигоді від використання.

В.2.3. Можливість перевірки результату. Повинні існувати методи визначення абсолютно кращого варіанту розв'язання задачі. Це дозволить не тільки порівнювати алгоритми між собою, а й оцінювати їх абсолютну ефективність.

В.2.4. Складність рішення. Типові задачі не мають бути тривіальними. Досягнення абсолютно кращого результату повинно бути практично неможливим.

В.2.5. Можливість порівняння результатів. Постановка задачі повинна забезпечувати можливість визначення якості результатів різних її рішень. Іншими словами крім таких результатів, як «задачу вирішено» і «задачу не вирішено» має існувати також безліч проміжних градацій.

В.2.6. Репрезентативність. Чим більше різних економічних задач може бути зведене до тих же постановок, що і типова, тим краще (рис. 2.16).

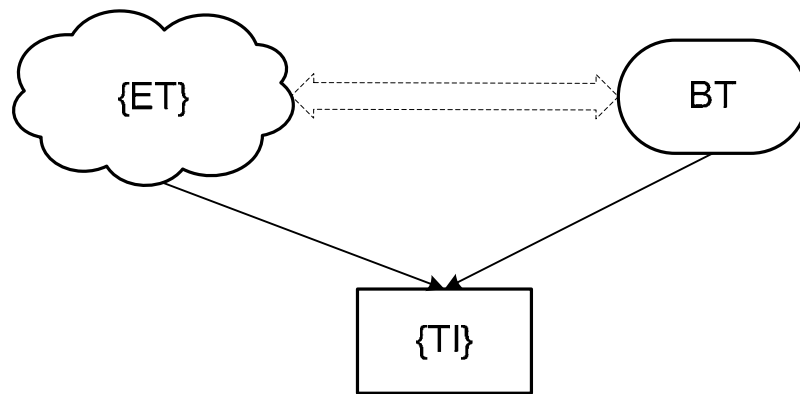


Рисунок 2.16 – Взаємовідносини множини економічних задач {ET} з типовими BT через множину постановок {TI}

Слід зазначити, що *de facto* в сфері інтелектуальних обчислень вже склався певний набір задач, які традиційно використовуються для перевірки роботи алгоритмів і демонстрації їх можливостей.

Так, довгий час однією з них була задача класифікації, в якій потрібно віднести рослину ірис до одного з трьох видів (*setosa*, *versicolour* або *virginica*) в залежності від довжини і ширини чашолистків і пелюсток. Аналіз відповідності вимогам В.2.1 - В.2.6 показує, що ця задача повністю відповідає лише вимогам В.2.3 і В.2.5, а також частково – В.2.6. Таким чином, підстав використовувати її в якості типової немає. Це розуміють і виробники програмного забезпечення, тому в демонстраційних прикладах сучасних систем нейромережевого моделювання вона майже не зустрічається.

Для аналізу ефективності рішення NP-повних задач часто використовується задача комівояжера – одна з найвідоміших задач комбінаторної оптимізації [96, 121]. Суть її зводиться до відшукування найкоротшого шляху обходу заданих міст, з подальшим поверненням в початок маршруту. Вперше вона була поставлена в XIX столітті, а в якості математичної проблеми отримала сучасне формулювання і назву на початку 1930-х років. Починаючи з 1950-х років, ведеться систематичний пошук аналітичних рішень цієї задачі, що дозволило отримати методи і алгоритми пошуку найкоротших шляхів для досить великої кількості вузлів. Недоліком цих методів є досить вузька спеціалізація, проте їх застосування дозволяє виконати вимогу В.2.3 до типових завдань. Інші вимоги також виконуються. Дані легко можуть бути згенеровані (В.2.1). Найкоротший маршрут є зазвичай і найбільш вигідним економічно (В.2.2). Знаходження найкоротшого шляху є дійсно складною задачею (В.2.4). Відношення знайденого шляху до найкоротшому дозволяє однозначно визначити ефективність різних алгоритмів (В.2.5) До комбінаторної оптимізації в економіці можна звести велику кількість практичних завдань (В.2.6).

Ще однією задачею, яка має велике значення для порівняння алгоритмів реалізації інтелектуальних обчислень, є задача біржового спекулянта. Вона може бути використана в якості типової для задач класифікації, регресії і кластеризації на слабозв'язаних даних [161]. Суть задачі зводиться до отримання прибутку від різниці курсів купівлі та продажу валюти, або інших

фінансових інструментів на біржі. Розглянемо відповідність задачі раніше сформульованим критеріям. Результати біржових торгів є відкритими і загальнодоступними, що забезпечує виконання вимоги В.2.1. Отримані результати прямо виражаються через прибуток, що забезпечує прозорість їх економічної інтерпретації (В.2.2). Абсолютно кращий варіант рішення відповідає точному прогнозу розвитку подій, що відомо з аналізу даних передісторії. Це забезпечує виконання вимог з перевірки результату (В.2.3). Досягнення абсолютно кращого результату на практиці неможливо, оскільки на розвиток подій впливає безліч факторів, які лише побічно проявляються у вхідних даних (В.2.4). Результати роботи різних алгоритмів можна зіставити за розміром отриманого прибутку (В.2.5). Що стосується репрезентативності результатів (В.2.6), то вище вже зазначалося, що дана задача може розглядатися в різних постановках (зокрема – регресії, класифікації, кластеризації), до яких можна звести велику кількість різних економічних задач з аналізу слабозв'язаних даних.

Використання системи типових задач в поєднанні з розробленими раніше методами забезпечення порівнянності результатів [195, 124, 54 и др.] дозволяє підвищити обґрунтованість вибору програмних засобів і алгоритмів навчання та знизити витрати на створення ІСПР. При цьому система типових задач може формуватися поступово, на підставі узагальнення накопиченого досвіду в створенні ІСПР.

Ефективність машинного навчання. Термін «навчання» стосовно до систем штучного інтелекту (комп'ютерних програм) визначено в класичній монографії Т.Мітчелла наступним чином [46, с.2]:

Кажуть, що комп'ютерна програма навчається при вирішенні якоїсь задачі з класу T , якщо її продуктивність, згідно метриці P , поліпшується при накопиченні досвіду E .

При цьому P , T і E можуть мати різні значення для різних задач.

Отже, момент припинення росту продуктивності системи є очевидним маркером закінчення процесу навчання. Однак на практиці виявлення цього моменту часто є нетривіальною задачею, зважаючи на складність визначення продуктивності інтелектуальної системи для різних класів задач. Крім того, поняття «навчання», відповідно до визначення, даного вище, можна розглядати у вузькому і в широкому сенсах.

Визначення 2.2.

У вузькому сенсі під навчанням будемо розуміти настройку параметрів програми спеціалізованим навчальним алгоритмом (наприклад, Back Propagation для ШНМ, C4.5 для дерев рішень, оператори селекції генетичних алгоритмів).

Визначення 2.3.

У широкому сенсі під навчанням будемо розуміти настройку самих навчальних алгоритмів, а також зокрема визначення структури системи інтелектуальних обчислень (наприклад - кількість нейронів в прихованих шарах ШНМ, кількість розгалужень в деревах рішень, розмір популяції в генетичних алгоритмах).

Питання визначення оцінки продуктивності системи штучного інтелекту найприродніше вирішується для задач аналізу даних, що відносяться до прогностичної групи (класифікація і регресія). Вхідна вибірка даних в цьому випадку ділиться на навчальну та тестову, а якість навчання визначається за наступними критеріями (К.2.3.1 – К.2.3.3) [180, с. 96]

К.2.3.1. Досягнення незначної помилки у розпізнаванні навчальної множини;

К.2.3.2. Досягнення незначної помилки у розпізнаванні тестової множини;

К.2.3.3. Досягнення адекватної динаміки процесу навчання

Пояснимо критерій К.2.3.3. Аналіз динаміки навчання нейронної мережі дозволяє виявити момент перенавчання ШНМ, що є основною небезпекою роботи з малими вибірками даних, та уникнути його. На рис. 2.17 показано класичний вид графіка зміни середньої помилки ШНМ на навчальній та тестовій множині даних.

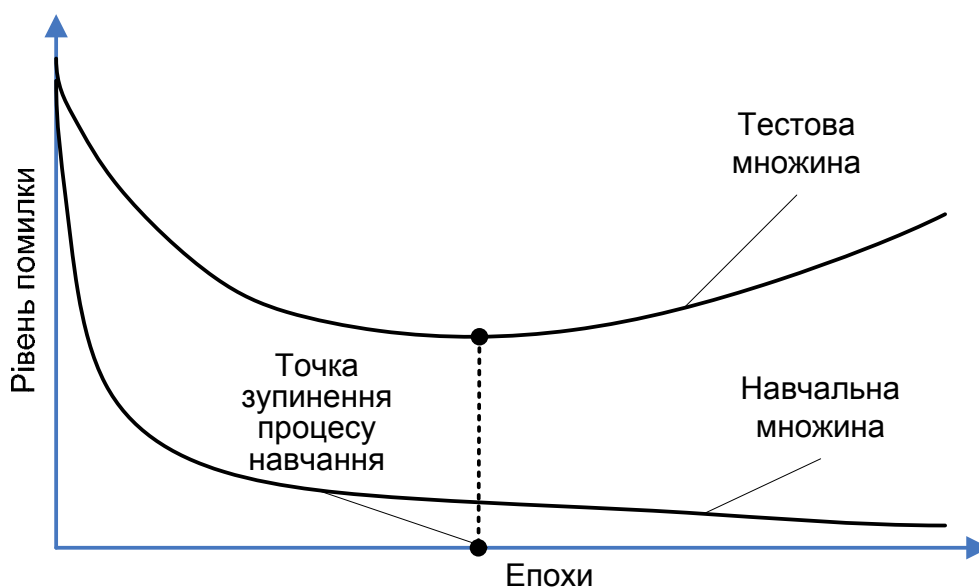


Рисунок 2.17 – Динаміка процесу навчання ШНМ

З аналізу рис. 2.17 видно, що якщо середня помилка на навчальній множині монотонно зменшується, то середня помилка ШНМ на тестовій множині проходить через екстремум, після чого починає збільшуватися. Точка цього екстремуму і є моментом, коли процес навчання необхідно зупинити. В іншому випадку ШНМ замість пошуку залежностей між вхідними та вихідними даними почне «запам'ятовувати» приклади з навчальної множини.

С. Хайкін вказує, що перехресна перевірка з використанням тестової множини і рання зупинка процесу навчання доцільні при виконанні наступної умови [232, с. 293]:

$$N < 30 W, \quad (2.25)$$

де W - кількість вільних параметрів нейронної мережі;

N - кількість прикладів в навчальній вибірці.

Тобто при малих обсягах навчальної вибірки перехресна перевірка необхідна.

Графік, показаний на рис. 2.17 також може служити зразком при аналізі процесу навчання, оскільки при явній суперечливості характеристик навчальної вибірки і архітектури ШНМ, його вигляд буде істотно відрізнятися від наведеного.

При аналізі ефективності навчання в задачах, в яких відсутня можливість розбиття вхідної вибірки на тестову і навчальну, застосування критеріїв К.2.3.1 і К.2.3.2 неможливо. Якщо при цьому кількість ітерацій в процесі навчання не обмежена, може бути задіяний критерій К.2.3.3, відповідно до якого маркером закінчення процесу навчання слід вважати припинення поліпшення значень функції продуктивності. Наприклад, для генетичних алгоритмів такої буде функція пристосованості, а для самоорганізаційних ШНМ із шаром Кохонена - функція помилки.

Відзначимо, що оцінка продуктивності інтелектуальної системи у вузькому сенсі за критерієм К.2.3.3 не є можливою для алгоритмів навчання з кінцевою кількістю ітерацій, наприклад для алгоритмів синтезу дерев рішень, алгоритмів кластеризації і тому подібних. Взагалі, непряма оцінка ефективності навчання за критеріями К.2.3.1 - К.2.3.3 для таких випадків не має сенсу. Підбір параметрів навчання таких алгоритмів здійснюється в комплексі з оцінкою ефективності результатів методами комбінаторної оптимізації.

Оцінка ефективності результатів аналізу даних. При оцінці результатів вирішення економічних задач головним критерієм ефективності вирішення є економічна вигода. Інакше кажучи, ефективність позначає здатність розробленої моделі виконувати поставлені практичні завдання.

Оцінка ефективності отриманої моделі може здійснюватися або в режимі реальної експлуатації, або на підставі даних про минулий стан системи. Режим реальної експлуатації дозволяє виявити *справжню* ефективність рішення, але ця перевірка обходиться дорожче і займає багато часу. Перевірка на підставі минулих даних може бути проведена набагато швидше і коштує дешевше, але вона дозволяє знайти тільки *очікувану* ефективність і її достовірність падає в умовах мінливого зовнішнього середовища. Таким чином, обидва методи мають як переваги, так і недоліки і на практиці застосовуються спільно.

Розглянемо класифікацію моделей аналізу даних з погляду на оцінку їх ефективності. Вона ґрунтується на класифікації, даної в [195, с. 565] де в межах даного класифікаційного признаку виділено два типи моделей – кількісні та описові. Класифікація, яка пропонується нижче, виділяє три типи моделей. Дано їх визначення.

Визначення 2.4.

Кількісними моделями першого типу будемо називати моделі, економічна ефективність яких однозначно визначається кількісними метриками, заснованими на різниці між результатами, передбаченими моделлю і фактичними даними. До цього типу, наприклад, відносяться моделі, які вирішують задачу прогнозування.

Визначення 2.5.

Кількісними моделями другого типу будемо називати моделі, економічна ефективність яких також залежить від достовірності передбачених подій, але абсолютна вартість правильних і помилкових прогнозів різниться. Серед моделей аналізу даних до цього типу зокрема відносяться ті, що вирішують задачі бінарної класифікації.

Визначення 2.6.

До *deskриптивного типу* будемо відносити моделі, результат роботи яких носить характер опису і для яких застосування формальних методів оцінки точності неможливо. Такими є, наприклад, моделі, що відносяться до класу аналізу зв'язків.

Для кожного з перелічених типів існують методи оцінки ефективності, що розрізняються областями застосування, складністю реалізації та якістю оцінок. Розглянемо деякі з них.

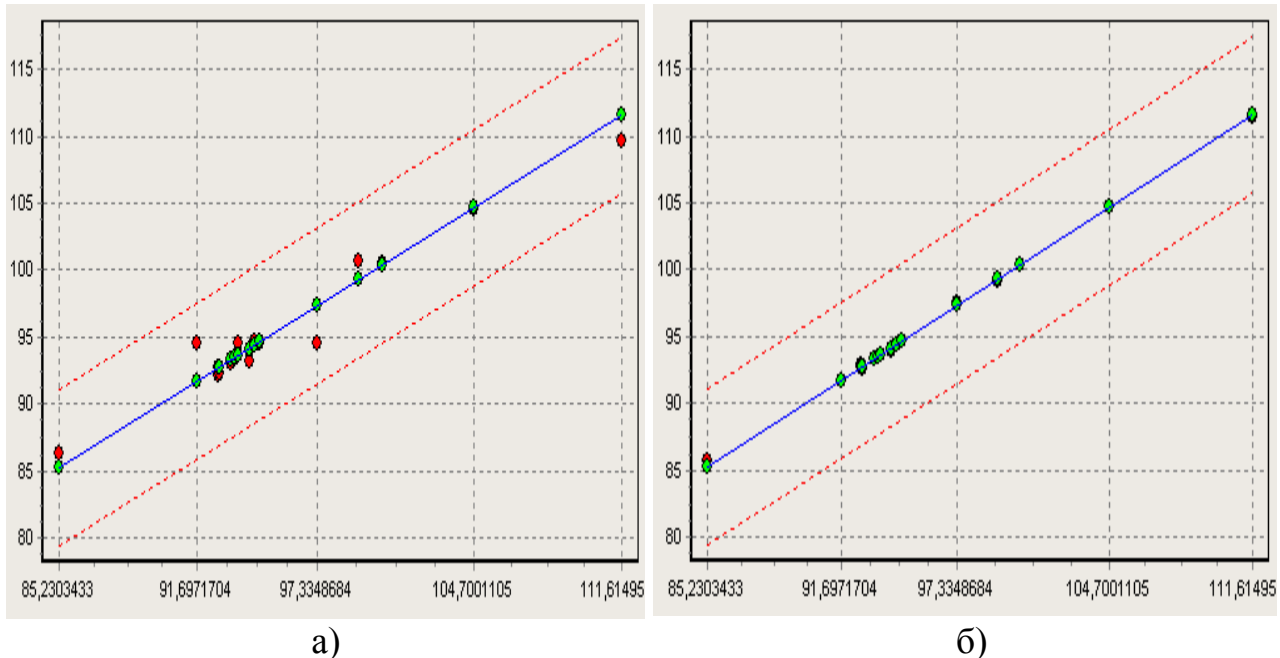
Ефективність *кількісних моделей першого типу* добре описується такою метрикою як середньоквадратична помилка прогнозування на навчальній і тестовій вибірках даних. У складних випадках, коли вибірка даних недостатньо репрезентативна, а крива динаміки процесу навчання істотно відрізняється від виду, показаного на рис. 2.17, можуть бути використані додаткові засоби візуальної оцінки ефективності прогнозування, зокрема, *діаграми розсіювання*, що дозволяють візуально оцінити ступінь і розподіл помилок прогнозування. На діаграмі у вигляді окремих точок відображаються вихідні значення кожного з прикладів навчальної вибірки, розраховані за допомогою моделі та ідеальні значення, що містяться в вибірці і розташовані на головній діагоналі. Вважається, що чим ближче розрахункові значення до ідеальних, тим менше помилка прогнозування. Однак така інтерпретація лише дублює показник середньоквадратичної помилки, тоді діаграми розсіювання можуть використовуватися і як інструмент для виявлення ефекту перенавчання [171].

Розглянемо діаграми, представлені на рис. 2.18. Вони отримані при аналізі ШНМ, навчених на вибірці даних, об'єм якої згідно (2.5) і (2.25) можна трактувати, як недостатній. Це викликає ефект перенавчання, в результаті чого нейронна мережа «запам'ятовує» вхідну вибірку даних і демонструє на ній дуже хороші результати прогнозування. На реальних же даних прогноз часто виявляється помилковим.

Для подолання ефекту перенавчання необхідно зменшити кількість вільних параметрів нейронної мережі до мінімуму, що забезпечує припустиме рішення задачі. Використання традиційного способу – розбиття вхідної вибірки на навчальну і тестову, з подальшим контролем помилки за тестовою вибіркою (рис. 2.17) в даному випадку не є достатньо ефективним з огляду на занадто малий розмір вхідної вибірки. У той же час, аналіз діаграм розсіювання різних моделей дозволяє отримати непряму інформацію про їх узагальнюючі здібності.

Так, діаграма, показана на рис. 2.18 а) відображає результати, отримані за допомогою ШНМ, що містить в прихованому шарі 2 нейрона. Діаграма, показана на рис. 2.18 б) відображає результати ШНМ, що містить в

прихованому шарі 3 нейрона. Незважаючи на те, що на діаграмі з рис. 2.18 б) показано практично ідеальне розпізнавання вхідних прикладів, в даних умовах це свідчить про перенавчання ШНМ, тому для подальшого використання доцільно вибрати мережу з меншою кількістю вільних параметрів, діаграма розсіювання якої наведена на рис. 2.18 а).



а) нормально навчена модель б) перенавчена модель
Рисунок 2.18 – Діаграми розсіювання моделей прогнозування

Ефективність кількісних моделей другого типу досліджується за допомогою інструментів, що дозволяють інтерпретувати результати інтелектуальних обчислень з погляду економічного ефекту і з урахуванням особливостей конкретної задачі. Суть цих інструментів зазвичай зводиться до розрахункових моделей, а також візуальних методів аналізу, серед яких можна виділити матриці спряженості, Lift-діаграми, ROC-криві і їм подібні [195]. Розглянемо деякі з них.

Розглянемо метрики, які використовуються для формальної оцінки точності бінарного класифікатора [54].

Нехай на розглянутій вибірці

TP – кількість правильно розпізнаних позитивних результатів;

TN – кількість правильно розпізнаних негативних результатів;

FP – кількість негативних результатів, розпізнаних як позитивні;

FN – кількість позитивних результатів, розпізнаних як негативні.

Тоді основні характеристики бінарного класифікатора можна описати наступними виразами.

Точність класифікатора показує, скільки з передбачених позитивних результатів виявилися дійсно позитивними [54]:

$$Prec = \frac{TP}{TP + FP}. \quad (2.26)$$

Повнота класифікатора показує, скільки із загальної кількості позитивних результатів було передбачене правильно [54]:

$$Rec = \frac{TP}{TP + FN}. \quad (2.27)$$

У практичних задачах одна з характеристик (*Prec* чи *Rec*) може виявитися важливіше за іншу. Наприклад, в задачах ранньої медичної діагностики важливіше повнота, а в разі класифікації потенційних банківських позичальників – точність.

Для того щоб дати загальну оцінку по точності і по повноті використовується *F-міра* [54]:

$$F_{mes} = \frac{1}{\alpha \frac{1}{Prec} + (1 - \alpha) \frac{1}{Rec}}, \quad \alpha \in [0,1], \quad (2.28)$$

де коефіцієнт α задає відношення ваг точності і повноти.

Існують і інші формули для формальної оцінки бінарних класифікаторів. Але частіше більш наочним і зручним є візуальний аналіз. Одним з його інструментів є матриця спряженості, приклад якої наведено в табл. 2.7 [153].

Таблиця 2.7 – Приклад матриці спряженості при аналізі банківських позичальників

	Класифіковано		
Фактично	Поганий	Добрий	Всього
Поганий	34	1	35
Добрий	3	111	114
Всього	37	112	149

Матриця спряженості відображає кількість правильно і неправильно класифікованих зразків із вхідної вибірки та дозволяє контролювати результати навчання в разі асиметрії ціни помилок класифікації першого і другого роду. До *помилки першого роду* відноситься класифікація поганих результатів як хороших, До *помилки другого роду* - класифікація хороших результатів як поганих [117].

Ціна помилок першого і другого роду може мати суттєві відмінності. Так при аналізі кредитоспроможності клієнта банк понесе істотно більші збитки, якщо поганого клієнта прийме за хорошого, ніж навпаки. Тобто в цьому випадку вартість помилок першого роду значно вище, ніж помилок другого роду.

Найбільш складним завданням з оцінки ефективності результатів аналізу даних є оцінка результатів *deskriptivних моделей*. Застосування апріорних методів оцінки економічного ефекту для них в більшості випадків неможливо, тому всі використовувані прийоми носять непрямий характер.

До таких прийомів відносяться принципи «бритва Оккама» та «мінімальна довжина опису» засновані на доказі теореми про те, що з кількох моделей, що описують дані з однаковою точністю, кращою є більш коротка [46]. Хоча в оригінальному дослідженні Т. Мітчелл пише про дерева прийняття рішень, дані принципи можуть бути поширені і на інші інструменти інтелектуальних обчислень, які вирішують схожі задачі. Зокрема рішення по діаграмах розсіювання, наведеним вище на рис. 2.18, також повністю відповідає їм.

Оцінка ефективності результатів обробки даних. В рамках ІСПР задачі обробки даних переважно носять забезпечувальний характер, тобто вирішуються для поліпшення результатів подальшого аналізу даних. В цьому випадку ефективність застосування k -го методу обробки даних щодо h -го визначається, як:

$$e_{dp}^{k/h} = \frac{e_s^k}{e_s^h}, \quad (2.29)$$

де e_s^k и e_s^h – ефективність всієї ІСПР, або тієї її частини, для якої вона може бути розрахована, за умови, що інші методи, складові ІСПР незмінні, за винятком внутрішньої структури, створюваної в результаті навчання.

Оцінки, отримані за допомогою виразу (2.29) є відносними, тоді як на практиці зручніше мати справу з абсолютними. Для переходу до абсолютних оцінок можна використовувати концепцію «наївного», або «марного» методу обробки, застосування якого гарантовано не покращує структуру вихідних даних. Для задач обробки даних *зі зміною порядку елементів* таким є випадковий вибір даних з вхідного масиву. Для задач обробки даних *без зміни порядку елементів* дані на виході «наївного» методу будуть збігатися з даними на його вході, тобто фактично обробка буде відсутня.

По відношенню до «наївного» будь-який метод, що дає по формулі (2.29) результат більше 1 може вважатися ефективним, а при результаті менше 1 - неефективним. При цьому виконується *правило транзитивності оцінок*, яке впливає з наступного твердження.

Твердження 2.3

Якщо метод k_1 краще «наївного» в x раз, а метод k_2 краще «наївного» в y раз, то метод k_1 краще k_2 в $\frac{x}{y}$ раз.

Доведення твердження 2.3

Припустимо ефективність «наївного» методу $e_s^0 = 1$.

Тоді за умовами твердження, та за формулою (2.29), ефективність методу

$k_1: e_s^{k_1} = x$; ефективність методу $k_2: e_s^{k_2} = y$.

$$\text{Тоді } e_{dp}^{k_1/k_2} = \frac{e_s^{k_1}}{e_s^0} \div \frac{e_s^{k_2}}{e_s^0} = \frac{e_s^{k_1} \cdot e_s^0}{e_s^0 \cdot e_s^{k_2}} = \frac{e_s^{k_1}}{e_s^{k_2}} = \frac{x}{y}$$

Твердження 2.3 доведено.

Використання правила транзитивності оцінок дозволяє істотно спростити процедуру оцінки методів обробки даних.

Існують також економічні задачі, вирішення яких безпосередньо зводиться до вирішення задачі обробки даних. В цьому випадку процедура оцінки ефективності передбачає визначення безпосереднього економічного результату від застосування методу. Серед них є задача ранжирування, до якої зводиться широке коло економічних задач, зокрема:

- ранжирування результатів обробки пошукових запитів за релевантністю;
- ранжирування клієнтів по ймовірності відгуку, в тому числі:
 - формування списку розсилки реклами;
 - формування списку боржників для робіт по стягненню боргу;
 - формування кредитних скорингових карт;
- ранжирування місць для відкриття філій і представництв торгових і фінансових установ за критерієм найбільшої відвідуваності клієнтами цільових груп;
- ранжирування методів навчання нейромережевих моделей за критерієм найбільшої ефективності роботи навченої моделі;
- ранжирування методів технічного аналізу валютних і фондових ринків за критерієм ефективності для побудови інтегральних оцінок;
- ранжирування змінних в порядку їх значимості для використання в інтегральних моделях (зокрема - нейромережевих).

Не зупиняючись на інструментальних засобах побудови моделей ранжирування, які можуть мати різну природу, в тому числі описову, розглянемо категорію ефективності ранжирування і методи, які можуть бути використані для її оцінки.

Введемо наступні позначення:

A_i – вектор параметрів, що характеризують i -й об'єкт, який підлягає ранжируванню. При цьому набір параметрів однаковий для всіх таких об'єктів.

$\{A\} = \{A_1, A_2, \dots, A_n\}$ – множина вхідних даних;

rp_i – рангова ознака i -го об'єкту, яка визначена за допомогою моделі ранжирування на підставі A_i і обумовлює позицію об'єкта в ранжируваному списку;

$RP = \{rp_1, rp_2, \dots, rp_n\}$ – вектор рангових ознак, визначених моделлю:

$$RP = M_r(A), \quad (2.30)$$

де M_r – модель ранжирування.

trp_i – істинна рангова ознака – зовнішня характеристика, яка в загальному випадку не може бути знайдена з A_i , але визначається *post factum* на підставі

додаткових даних;

$TRP = \{trp_1, trp_2, \dots, trp_n\}$ – вектор дійсних рангових ознак;

r_i – ранг (місце в рейтингу) i -го об'єкту, визначений на підставі rp_i . $r_i \in \mathbb{N}$;

$R = \{r_1, r_2, \dots, r_n\}$ – вектор рангів, визначених моделлю:

$$R = rang(RP); \quad (2.31)$$

tr_i – істинний ранг i -го об'єкта, визначений на підставі trp_i . $tr_i \in \mathbb{N}$;

$TR = \{tr_1, tr_2, \dots, tr_n\}$ – вектор істинних рангів:

$$TR = rang(TRP); \quad (2.32)$$

Таким чином, визначення ефективності ранжирування може бути зроблено виходячи з вектору відстаней між істинними рангами об'єктів і рангами, визначеними моделлю. $\Delta(R, TR) = \{d_1, d_2, \dots, d_n\}$.

Похибку ранжирування моделі Mr , тобто величину, яка зворотно пропорційна її ефективності в цьому випадку можна визначити так:

$$eMr = \sum_{i=1}^n |d_i|. \quad (2.33)$$

Для визначення відстані між дійсними рангами і рангами, визначеними моделлю (d_i) в найпростішому випадку можна скористатися виразом:

$$d_i = r_i - tr_i. \quad (2.34)$$

На практиці при обчисленні d_i и eMr необхідно враховувати додаткові фактори, які визначаються економічною сутністю розв'язуваної задачі:

1. Задача може допускати існування об'єктів з однаковими рангами, тобто $tr_i = tr_j$ для деяких $i, j = 1..n$. Це означає, що об'єкт може зайняти будь-яке з деякої кількості місць і це не вплине на ефективність вирішення. Обчислення d_i в такій задачі за допомогою формули (2.34) призведе до помилкового завищення похибки ранжирування, отже, при розробці алгоритму обчислення d_i необхідно забезпечити врахування даного чинника.

2. Значимість правильного ранжирування об'єктів на початку і в кінці списку може бути різною. Так, у багатьох задачах (пошуку інформації, ранжирування об'єктів в порядку убуття переваг і тому подібні) найбільшу вагу має правильне розташування об'єктів у верхній частині списку, тоді як за її межами допускаються досить сильні відхилення знайденого рейтингу від дійсного. Для урахування фактора значущості, вираз (2.33) можна доповнити:

$$eMr = \sum_{i=1}^n f(i) |d_i|, \quad (2.35)$$

де $f(i)$ – функція значущості i -го місця в рейтингу.

Крім суто розрахункових методів, для аналізу ефективності ранжирування можуть бути застосовані і розрахунково-графічні методи, зокрема Lift-криві і їх різновиди (Profit-криві, Gain-діаграми та інші) [195].

Lift-крива формується на основі ліфт-фактора, який був початково визначений при вирішенні задачі оптимізації масової розсилки як показник, що відображає збільшення числа відгуків щодо числа дій (поштових відправлень). Lift-крива будується наступним чином: по горизонтальній осі відкладається розмір вибірки впорядкованої по спаду показника rp_i , який відображає ймовірність настання позитивного результату, згідно з аналізованою моделлю. За вертикаллю фіксується кумулятивне число позитивних результатів в кожній підвибірці (ліфт). Оскільки істинна рангова ознака trp_i в даному випадку є бінарної величиною, яка приймає значення 1 в разі позитивного результату i та значення 0 у разі негативного, вираз для розрахунку ліфта можна записати в такий спосіб:

$$lft(i) = \sum_{x=1}^i trp_x. \quad (2.36)$$

Але класична Lift-крива може бути використана для оцінки ефективності моделей ранжирування тільки при бінарній природі істинних рангових ознак trp_i .

Як приклад використання Lift-кривої розглянемо задачу ранжирування прострочених кредитних справ в колекторському скорингу, в яких об'єкти необхідно розташувати за рівнем зменшення ймовірності відновлення позичальником платежу по кредиту. Для аналізу значущості використовується вибірка тестових даних, в яких істинний ранг кожної кредитної справи вже відомий. Так, аналізована вибірка містить дані по 500 позичальникам, з яких в результаті проведених банком заходів 83 позичальники відновили внесення кредитних платежів. Вектор вхідних даних має наступну структуру: {*стать позичальника; вік; сума прострочення; сума платежу; відношення прострочення/платіж; період прострочення; сума кредиту*}.

Припустимо, що для ранжирування використовується три моделі:

Перша модель є емпіричною і використовується в багатьох кредитних підрозділах для ранжирування даних про позичальників. В якості рангових ознак rp_i в ній обрано кількість днів, що минуло з моменту останнього платежу. Дійсність моделі підтверджують статистичні дослідження, відповідно до яких чим більше період прострочення, тим менше ймовірність відновлення оплати.

Другу модель засновано на визначенні зв'язку між вхідними та результуючим показником за допомогою логістичної регресійної моделі. Ця модель дозволяє забезпечити урахування більшої кількості характеристик

позичальника, ніж попередня, а також визначити їх значимість і вплив на ймовірність відновлення платежу.

Третя модель передбачає відсутність ранжирування, тобто випадковий вибір позичальника. Таку модель назвемо «марною». Вона служить візуальним орієнтиром для аналізу та зіставлення ефективності інших моделей та завжди присутня на діаграмі з Lift-кривими.

Графіки Lift-кривих перелічених моделей наведено на рис. 2.19.

Практичний сенс зіставлення ефективності за допомогою Lift-кривих полягає у визначенні моделі, що дозволяє зробити найменшу кількість дій, необхідних для досягнення певного результату.

В якості загального критерію оцінки моделі використовується площа під кривою, що виражена у відсотках. Для «марної» моделі цей показник завжди дорівнює 50%. Для моделі логістичної регресії з рис.2.19, площа під кривою – 63,3%. Для наведеної там же емпіричної моделі ранжирування – 55,1%. Таким чином, з розглянутих моделей за результатами зіставлення більш ефективною для ранжирування позичальників є модель логістичної регресії.

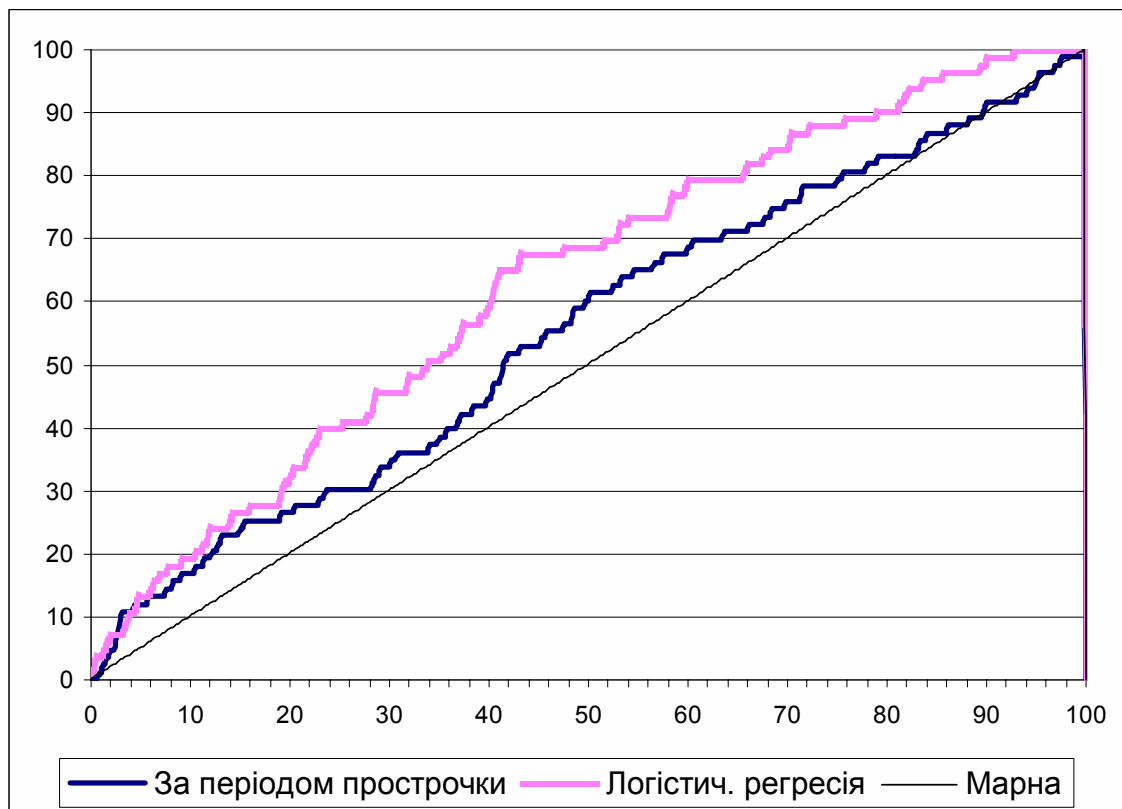


Рисунок 2.19 – Lift-криві різних моделей ранжирування позичальників в колекторському скорингу

Іншим критерієм оцінки ефективності вирішення задач ранжирування є частка відгуків, одержуваних при здійсненні певної кількості дій. Так, з аналізу графіків на рис. 2.19 можна побачити, що для отримання 50% позитивних відгуків в «марній» моделі необхідно обробити 50% кредитних справ, з використанням моделі логістичної регресії – 31% кредитних справ, а з

використанням моделі оцінки ймовірності погашення за періодом прострочення – 41% справ. Очевидно, що в даному випадку також доцільно вибрати саме модель логістичної регресії. При цьому слід зазначити, що емпірична модель демонструє хороші результати при обробці малої кількості заявок. Отже, можливість дослідження не тільки загальної ефективності моделей, а й ефективності їх на окремих ділянках діапазону ранжирування, є безумовною перевагою розрахунково-графічних методів оцінки ефективності вирішення задач ранжирування.

Але недоліком класичних Lift-кривих є неможливість їх використання для аналізу ефективності моделей ранжирування в тому випадку, якщо елементи множини дійсних рангових ознак trp_i мають не бінарну природу. Однак, даний метод, може бути вдосконалений для аналізу ефективності при довільній природі дійсних рангових ознак [157]. Вираз для розрахунку ліфта при цьому:

$$lft(i) = \sum_{x=1}^i \{r \leq tr_x\}, \quad (2.37)$$

де $\{r \leq tr_x\} = 1$ якщо істинно і 0, якщо помилково.

Метод побудови *кривої ефективності ранжирування*, заснований на функції (2.37), дозволяє забезпечити пріоритетну значимість правильного ранжирування об'єктів на початку списку. Дійсно, якщо за основний критерій ефективності моделі ранжирування прийняти площу під кривою (в частках від загальної площі графіка):

$$S_{lft} = \frac{\sum_{i=1}^n lft(i)}{n^2}, \quad (2.38)$$

то правильне розташування об'єкта на першому місці в ранзі збільшить загальну площу під кривою на величину $S_{lft} = \frac{n}{n^2}$, тоді як правильне

розташування об'єкта на останньому місці в ранзі тільки на величину $S_{lft} = \frac{1}{n^2}$.

Таким чином, значимість місць об'єктів в ранзі в запропонованому методі оцінки ефективності (2.37) убуває в арифметичній прогресії.

На рис. 2.20 наведено два крайніх випадки кривих ефективності ранжирування для марної та ідеальної моделей.

Діагональна пряма лінія ($S_{lft} = 0.5$) відповідає ідеальній моделі ранжирування, яка розташовує всі об'єкти на вірних місцях. Увігнута лінія під діагоналлю ($S_{lft} = 0.33$) відповідає марній моделі, в якій вибір місць в ранзі здійснюється випадково.

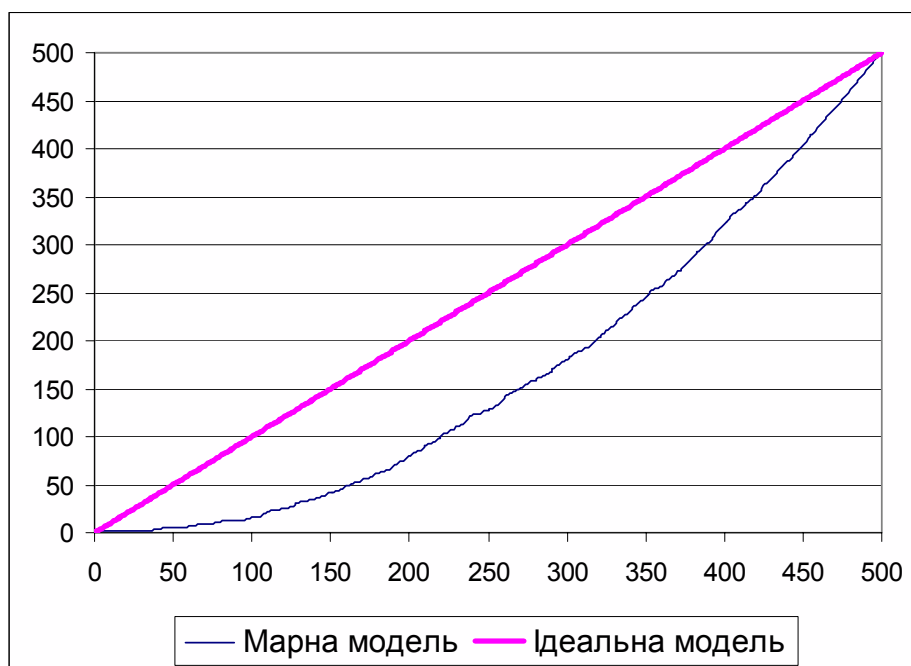


Рисунок 2.20 – Криві ефективності ранжирування марної та ідеальної моделі

Візуальний аналіз графіків на рис. 2.20 та їх зіставлення із графіками на рис. 2.19 показує, що і зовнішній вигляд і інтерпретація побудованих кривих ефективності істотно відрізняються від Lift-кривих, що вдосконалює методи графічного аналізу ефективності вирішення задач ранжирування.

Таким чином, проведення дослідження дозволило розвинути методологію моделювання штучних нейронних мереж для вирішення економічних задач за напрямками вибору інструментів моделювання, визначення архітектури штучних нейронних мереж, вдосконалення методів роботи із вхідними даними, дослідження ефективності вирішення економічних задач методами нейромережевого моделювання.

Розроблено низку моделей для здійснення обґрунтованого вибору інструментальних засобів реалізації інтелектуальних обчислень. Для умов великій кількості аналізованих продуктів і критеріїв запропоновано нечітку модель оцінки та формалізовано процедури фазифікації вхідних даних і дефазифікації результатів.

Систематизовано та розширено підходи до оцінки результатів інтелектуальних обчислень. Запропоновано та деталізовано концепцію типових задач. Запропоновано розрахункові та розрахунково-графічні методи вирішення проблеми оцінки результатів обробки даних. Перевагою розрахунково-графічних методів є наочність отриманих результатів. Перевагою суто розрахункових – кращі можливості з формалізації вибору оптимальних моделей.

РОЗДІЛ 3.

МОДЕЛІ ІННОВАЦІЙНИХ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ АНАЛІЗУ І ОБРОБКИ ДАНИХ В ЕКОНОМІЧНИХ СИСТЕМАХ

3.1 Нейромережеві моделі аналізу даних в економічних системах

Розглянемо особливості практичного вирішення економічних задач в різних базових постановках з використанням штучних нейронних мереж. В якості типових задач візьмемо задачу біржового спекулянта і задачу прогнозування банкрутств.

Рішення задачі біржового спекулянта. В р. 2 було показано, що задачу біржового спекулянта можна привести мінімум до трьох різних постановок – класифікації, регресії і кластеризації. Розглянемо реалізацію задачі в цих постановках.

Базовий набір вхідних даних, що описує стан валютного ринку, в загальному випадку містить інформацію про курсові відносини різних пар валют за часовими періодами. При цьому в кожному періоді виділяється курс на початок періоду o_i , на кінець періоду c_i , а також максимальне h_i та мінімальне l_i значення курсу за період.

Практичну реалізацію різних постановок задачі біржового спекулянта проведемо на прикладі аналізу щоденних даних про зміну курсу в валютній парі євро-фунт (EUR / GBP). Ця пара валют порівняно рідко буває об'єктом великих біржових спекуляцій і тому має низьку волатильність, внаслідок чого краще підходить для нейромережевого аналізу, ніж такі популярні в біржовій торгівлі пари валют, як євро – долар США (EUR / USD), або фунт-долар США (GBP / USD). Для аналізу взято дані за період з жовтня 2007 по грудень 2015 року. Для забезпечення урахування динаміки зміни курсів дані перетворені ковзаючим вікном із глибиною занурення 10 періодів. Тобто, для опису ринкової ситуації замість вектора $\{o_i, h_i, l_i, c_i\}$ використовується вектор $\{o_{i-10}, h_{i-10}, l_{i-10}, c_{i-10}, o_{i-9}, \dots, o_{i-1}, h_{i-1}, l_{i-1}, c_{i-1}, o_i, h_i, l_i, c_i\}$. Хоча це не єдиний можливий метод попередньої підготовки даних, для зіставлення ефективності вирішення задачі в різних постановках його застосування цілком виправдано. Всього вхідна вибірка даних налічує 2134 записи.

Виконання біржових операцій здійснюється за стандартною моделлю, яка передбачає відкриття позиції (покупку або продаж валюти) на початку прогнозованого часового інтервалу (дня) і закриття позиції (зворотну операцію) в кінці інтервалу.

Основним критерієм оцінки ефективності прогнозування в задачі біржового спекулянта є дохід, отриманий від операцій. При цьому дохід може бути розрахований як на всій вибірці вхідних даних, так і на її перевіірочній частині, тобто підмножині даних, яка не використовувалася для навчання мережі. При цьому перевіірочні дані можуть розташовуватися як в кінці вибірки, так і бути довільно розподілені по ній. Останній варіант дозволяє

зменшити вплив змін ринкових умов на отримані результати.

Розглянемо задачу біржового спекулянта, як задачу регресії. Результатом роботи моделі в цьому випадку є чисельний прогноз величини валютного курсу в наступний період, або прогноз величини відхилення валютного курсу в наступному періоді від поточного значення (див. рис. 2.5).

При вирішенні задачі в даній постановці вхідний вектор формують змінні $\{o_{i-10}, h_{i-10}, l_{i-10}, c_{i-10}, o_{i-9}, \dots, o_{i-1}, h_{i-1}, l_{i-1}, c_{i-1}, o_i\}$, а вихідний, тобто вектор значень, які повинна прогнозувати нейронна мережа, змінні $\{h_i, l_i, c_i\}$, або просто $\{c_i\}$.

Для вирішення задачі скористаємося можливостями аналітичної платформи Deductor. ШНМ яку створено для цього має структуру 43-8-3, тобто 43 нейрона у вхідному шарі, 8 нейронів в прихованому шарі і 3 нейрона у вихідному шарі. Кількість нейронів у прихованому шарі обрано з компромісних міркувань швидкості навчання і забезпечення здатності апроксимації мережею складних залежностей в біржових даних.

Розглянемо за якими критеріями можна оцінити якість навчання створеної ШНМ. Теорія нейронних мереж пропонує взяти за основний критерій якості навчання рівень помилки на навчальній і тестовій вибірках. Графік зміни цих параметрів в процесі навчання нейронної мережі показано на рис. 3.1.

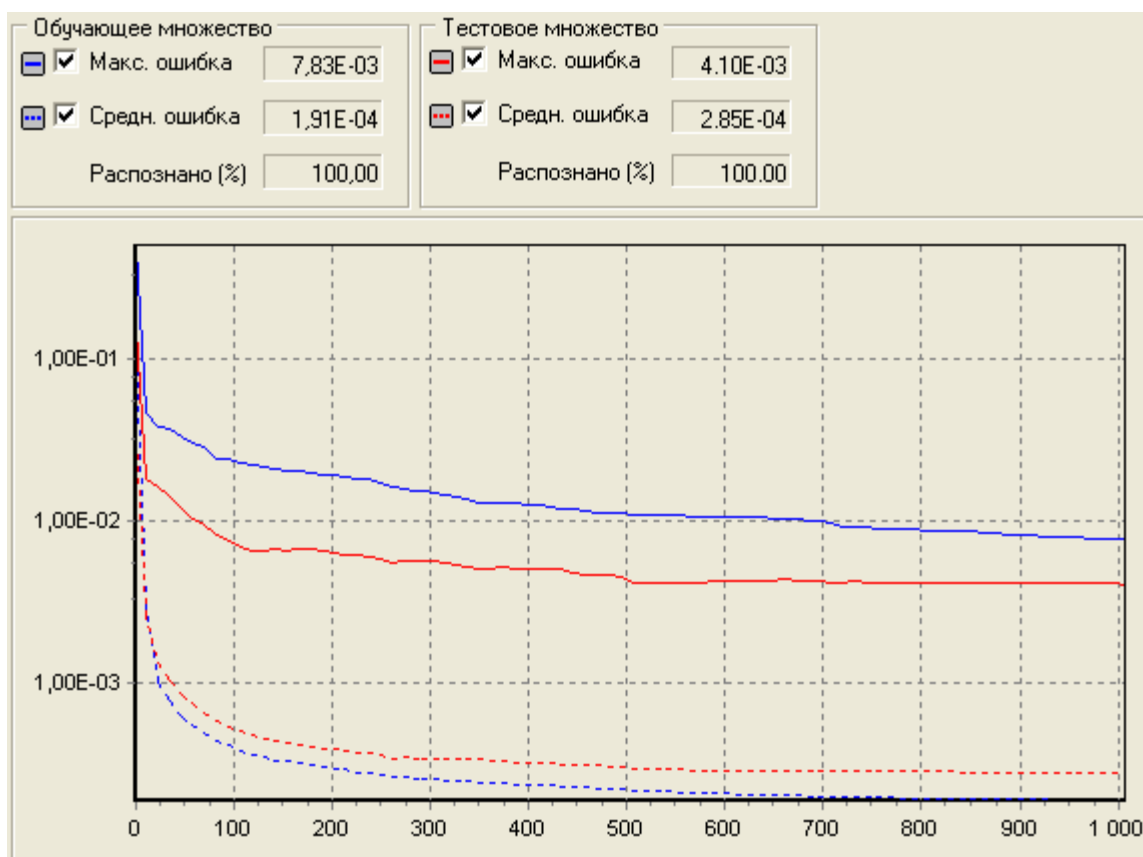


Рисунок 3.1 – Динаміка навчання ШНМ при вирішенні задачі біржового спекулянта, як задачі регресії

Як видно з аналізу рис. 3.1, в процесі навчання досягнуто достатньо низький рівень середньої помилки як на навчальній, так і на тестовій множинах.

Вид графіків, що відображають процес навчання також близький до класичного, що зазвичай свідчить про його нормальне протікання.

Проте, з практичної точки зору, набагато важливішими є показники прибутковості, що досягаються з використанням нейронної мережі. Відповідно до обраної моделі прийняття торгових рішень, покупка валюти (в нашому випадку EUR) здійснюється, якщо прогнозне значення курсу на кінець торгового періоду більше, ніж його значення на початок періоду, тобто якщо $c_i > o_i$. Продаж валюти здійснюється, якщо $c_i < o_i$. Якщо значення курсів рівні, транзакція не відбувається.

У табл. 3.1 наведено результати застосування отриманої моделі до аналізу реальних біржових даних з обраної пари валют.

Таблиця 3.1 – Результати біржових операцій по валютній парі EUR / GBP при регресійній постановці задачі нейромережевого моделювання

№ _{пп}	Спосіб розрахунку	Зроблено транзакцій					Результат торгівлі, пунктів	
		всього (при обсязі вибірки)	вдалих		невдалих		всього	на одну транзакцію
			абс.	%	абс.	%		
1	Розрахунок за всім обсягом даних (за навчальною і тестовою вибіркою)	2114 (2134)	1163	55	951	45	13369	6,26
2	Розрахунок тільки за тестовою вибіркою, яка розташована в кінці масиву даних	108 (108)	56	51,9	52	48,1	-321	-2,97
3	Розрахунок тільки за тестовою вибіркою, яка рівномірно розподілена по масиву даних	213 (214)	119	55,9	94	44,1	1177	5,53

Відповідно до традицій, прийнятих в біржових операціях, результат торгівлі розраховується в «пунктах», тобто одиницях, що відображають зміни курсу. Один пункт відповідає зміні курсу на 1 в останньому знаку. Існують також інші методи оцінки ефективності торгових систем [135], але в даному випадку їх використання не виправдане. Звернемо увагу на результат, отриманий при здійсненні операцій на тестовій вибірці, розташованій в кінці масиву даних. Негативний результат торгівлі, а також найменшу частку вдалих транзакцій в загальній їх кількості, в порівнянні з іншими методами розрахунку, можна пояснити істотною зміною ринкових умов, в порівнянні з

тими, які були враховані нейронною мережею при навчанні. Виходячи з цього, більш правильним уявляється спосіб, при якому значення, що використовуються для тестової вибірки, рівномірно розподілені по всій множині вихідних даних.

Додатковий аналіз ефективності роботи нейронної мережі може бути проведено за допомогою графіка зміни накопиченого результату торгівлі (рис. 3.2). Графік побудований за тестовою вибіркою, яка рівномірно розподілена в масиві даних.

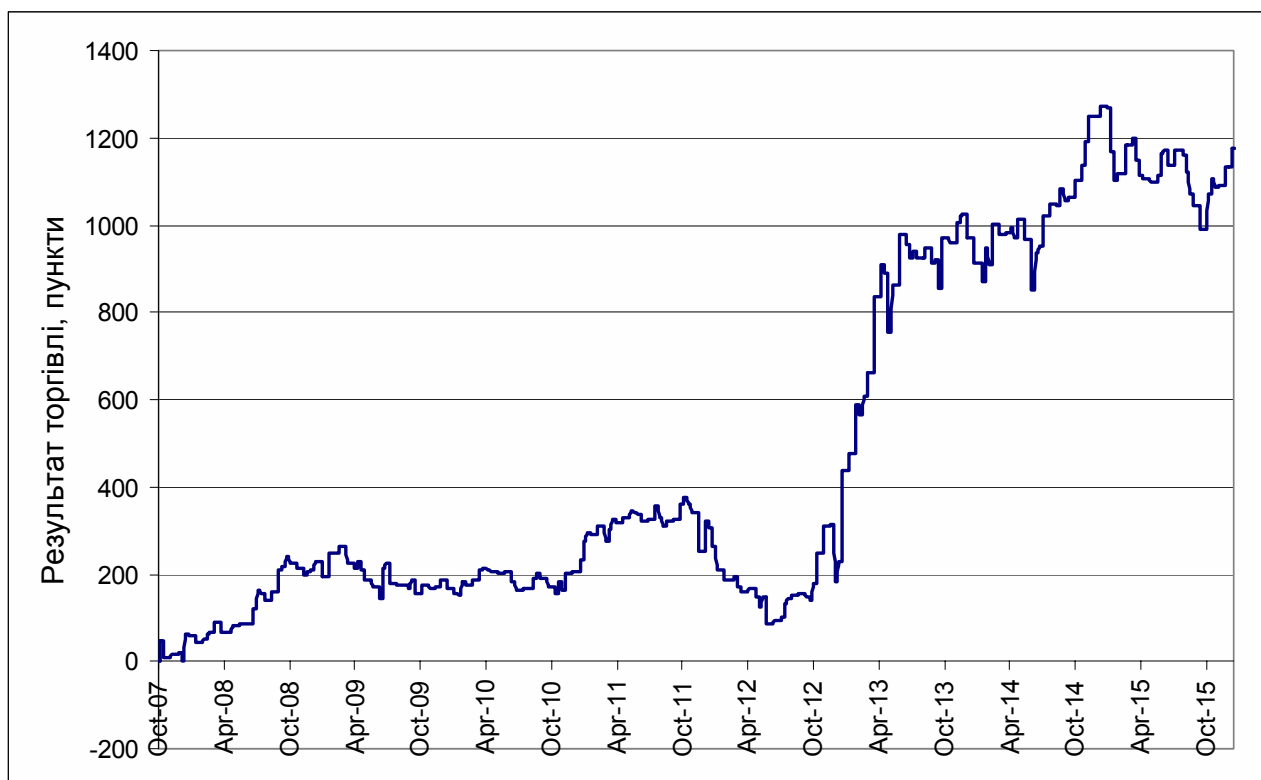


Рисунок 3.2 – Динаміка накопичення результатів торгівлі при вирішенні задачі біржового спекулянта, як задачі регресії

Як видно з аналізу рис. 3.2, виручка торгівельної системи стійко зростає та постійно знаходиться в зоні доходу. Це з позитивної сторони характеризує отриману модель

Розглянемо тепер рішення цієї ж задачі, як задачі класифікації. У цьому випадку від нейронної мережі потрібно класифікувати поточну біржову ситуацію, в залежності від дій, які рекомендується зробити валютному дилеру - продати валюту, купити валюту, або нічого не робити. Схематично це було показано на рис. 2.4.

Перед подачею біржових даних на вхід нейронної мережі, що вирішує задачу класифікації, необхідно провести їх додаткову обробку, суть якої полягає у визначенні очікуваного результату do_i по кожному вхідному вектору. На практиці при відносно невеликій кількості класів кількість виходів нейронної мережі може бути скорочено. Так в даному випадку можна обійтися двома або навіть одним виходом. При двох виходах для визначення очікуваних

результатів можна скористатися наступною формулою:

$$\begin{aligned} do_b_i &= \begin{cases} 1; & c_i > o_i \\ 0; & c_i \leq o_i \end{cases} \\ do_s_i &= \begin{cases} 1; & c_i < o_i \\ 0; & c_i \geq o_i \end{cases} \end{aligned} \quad (3.1)$$

де do_b_i – вихід, який відповідає рішення про купівлю валюти, do_s_i – вихід, який відповідає рішення про продаж валюти.

Якщо з яких-небудь причин ШНМ повинна мати тільки один вихід, для визначення do_i можна скористатися наступною формулою:

$$do_i = \begin{cases} 1; & c_i > o_i \\ -1; & c_i < o_i \\ 0; & c_i = o_i \end{cases} \quad (3.2)$$

При навчанні нейронної мережі для вирішення задачі класифікації параметри $\{do_b_i, do_s_i\}$ або $\{do_i\}$ складають вихідний набір змінних. Вхідний вектор при цьому, як і раніше, формують змінні $\{o_{i-10}, h_{i-10}, l_{i-10}, c_{i-10}, o_{i-9}, \dots, o_{i-1}, h_{i-1}, l_{i-1}, c_{i-1}, o_i\}$.

Результати застосування нейромережевої моделі класифікації для аналізу біржових даних з обраної парі валют наведено в табл. 3.2.

Таблиця 3.2 – Результати біржових операцій по валютній парі EUR / GBP при класифікаційній постановці задачі нейромережевого моделювання

№ _{тп}	Спосіб розрахунку	Зроблено транзакцій					Результат торгівлі, пунктів	
		всього (при обсязі вибірки)	вдалих		невдалих		всього	на одну транзакцію
			абс.	%	абс.	%		
1	Розрахунок за всім обсягом даних	243 (2134)	201	82.7	42	17.3	8837	36,37
2	Розрахунок за тестовою вибіркою, яка рівномірно розподілена по масиву даних	29 (214)	18	62,1	11	37,9	628	21,66

Як видно з табл. 3.2, розрахунок при розташуванні тестової вибірки в кінці масиву даних, не проводився, оскільки раніше було виявлено його недоліки. Аналіз табл. 3.2 показує, що загальний результат торгівлі в даному

випадку менш, ніж при вирішенні задачі в регресійній постановці. Однак показники, що характеризують якість прогнозування, у даній моделі набагато кращі, в порівнянні із даними табл. 3.1.

Розглянемо графік накопичення результатів торгівлі на тестовій вибірці при використанні класифікаційної моделі (рис. 3.3).

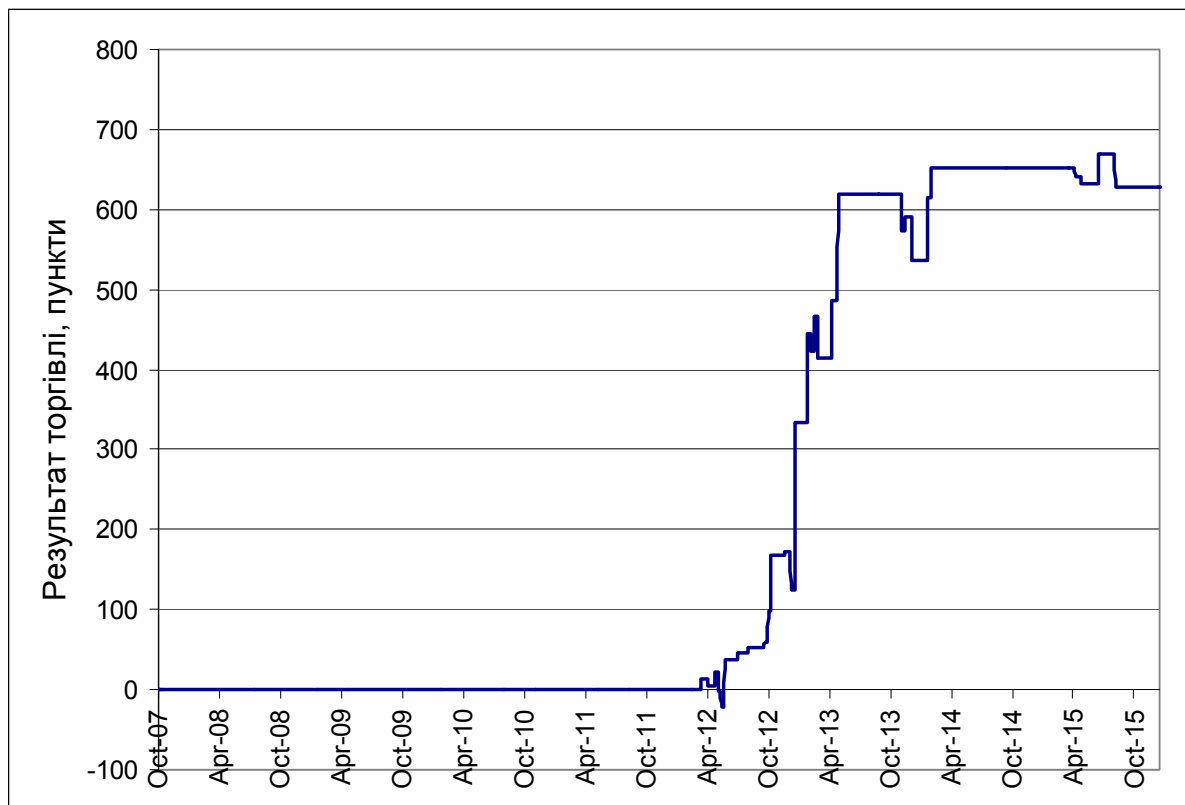


Рисунок 3.3 – Динаміка накопичення результатів торгівлі при вирішенні завдання біржового спекулянта, як завдання класифікації

Як видно з аналізу графіка на рис. 3.3, виручка торгової системи, як і в розглянутому раніше випадку (рис. 3.2) зростає достатньо впевнено. Однак кількість угод, укладених системою, набагато менше, ніж при вирішенні задачі у регресійній постановці.

Вирішення даної задачі як задачі угруповання об'єктів (кластеризації) здійснюється на підставі схеми з рис 2.6.

При обробці даних самоорганізаційною нейронною мережею вибірка даних, що використовуються при навчанні, складається тільки з векторів змінних $\{o_{i-10}, h_{i-10}, l_{i-10}, c_{i-10}, o_{i-9}, \dots, o_{i-1}, h_{i-1}, l_{i-1}, c_{i-1}, o_i\}$. Набір змінних, які для розглянутих раніше моделей були вихідними даними, при роботі з самоорганізаційними мережами використовується вже після закінчення навчання для інтерпретації змісту кластерів, виділених нейронною мережею.

В процесі моделювання на даній вибірці даних побудовано карту Кохонена з наступними параметрами: 49 комірок (розмір карти 7x7), 7 кластерів. Результати застосування нейромережевої моделі кластеризації до аналізу біржових даних з обраної пари валют наведено в табл. 3.3, а графік

накопичення прибутку в торгівельній системі показано на рис. 3.4.

Таблиця 3.3 – Результати біржових операцій по валютній парі EUR / GBP при постановці задачі нейромережевого моделювання як задачі кластеризації

№ _{пп}	Спосіб розрахунку	Зроблено транзакцій					Результат торгівлі, пунктів	
		всього (при обсязі вибірки)	вдалих		невдалих		всього	на одну транзакцію
			абс.	%	абс.	%		
1	Розрахунок за всім обсягом даних	2114 (2134)	1152	54,5	962	45,5	6555	3,10
2	Розрахунок за тестовою вибіркою, яка рівномірно розподілена по масиву даних	213 (214)	112	52,6	101	47,4	541	2,54

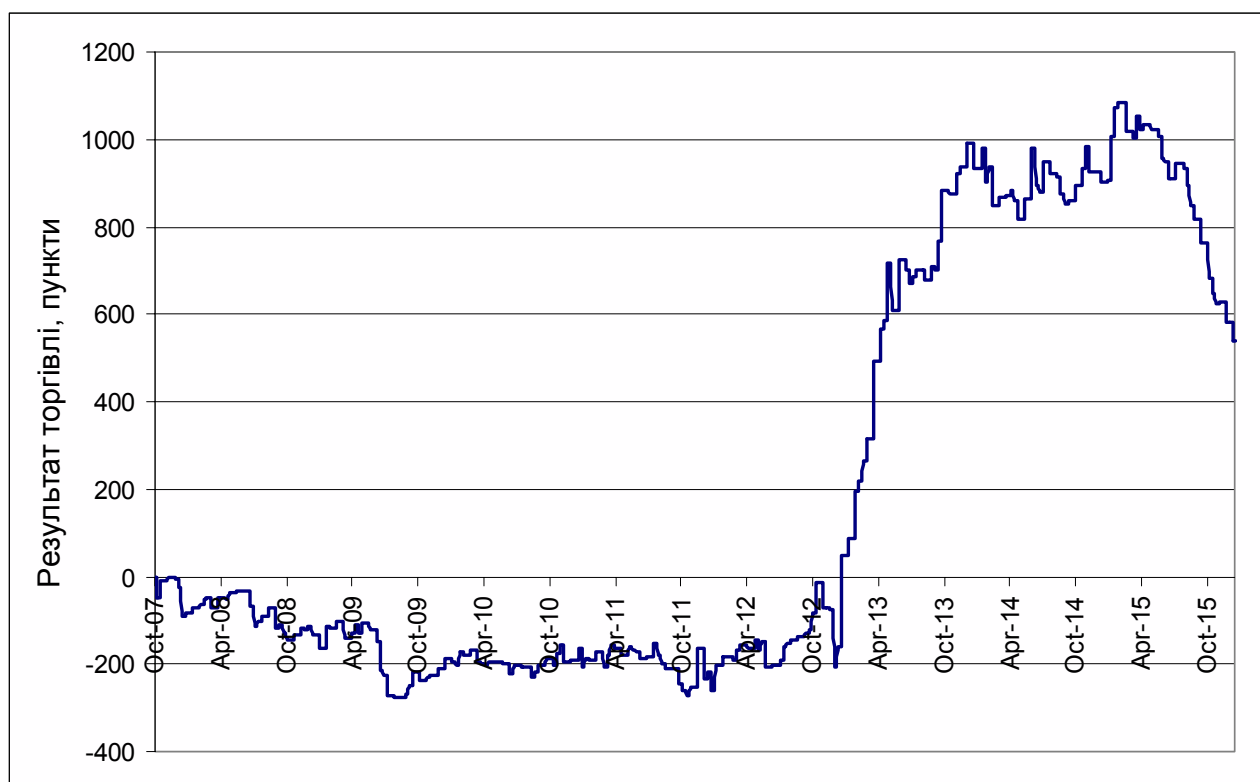


Рисунок 3.4 – Динаміка накопичення результатів торгівлі при вирішенні задачі біржового спекулянта, як задачі кластеризації

З аналізу табл. 3.3 видно, що рішення задачі біржового спекулянта в постановці кластеризації також дозволяє отримати позитивні результати торгівлі. Ефективність даної моделі за всіма параметрами виявилася менше, ніж у розглянутих вище моделей класифікації і регресії. Однак слід зазначити, що для отримання хороших результатів з самоорганізаційними нейронними

мережами потрібен більш досконалий підхід до підготовки вхідних даних та підбору параметрів ШНМ. Так, в [163] саме при використанні карт Кохонена було отримано найбільший економічний ефект.

З аналізу динаміки накопичення результатів торгівлі, представленої на рис. 3.4, можна зробити висновок, що потенційні можливості даної моделі досить високі. Так, різниця між мінімальними і максимальними значеннями торгового результату становить майже 1300 пунктів. Однак, алгоритм прийняття торгових рішень при роботі отримав багато дрібних збитків, що погіршило загальний результат.

В цілому, на підставі проведеного аналізу можна зробити висновок, що технології вирішення задачі біржового спекулянта у регресійній та класифікаційній постановках мало відрізняються одна від одної. Основна різниця виявляється в інтерпретації отриманих даних. При цьому, оскільки класифікаційна модель, забезпечує найкращі питомі показники, ніж регресійна, вона більше підходить для створення автоматичних торговельних систем. У той же час регресійна модель, як і кластеризаційна, забезпечує отримання прогнозу в кожному біржовому періоді, тому добре підходить для створення автоматизованих систем підтримки прийняття рішень.

Відносно низькі питомі показники прибутковості було отримано в результаті рішення задачі біржового спекулянта у регресійній і кластеризаційній постановках. Значною мірою це обумовлено застосуванням спрощеної методики прийняття торгових рішень, яка не враховує силу прогнозованих рухів валютного курсу. Однак якщо ці моделі застосовуються для створення систем підтримки прийняття рішень, то цей недолік не має принципового значення, оскільки досвідчений трейдер здатний фільтрувати інформацію про малоімовірні зміни курсу.

Рішення задачі прогнозування банкрутств. Розглянемо іншу задачу, яка в різних варіаціях також часто зустрічається при аналізі економічних систем. Вона пов'язана із прогнозуванням економічної неспроможності, або банкрутства контрагентів. [172, 173].

Актуальність задачі обумовлено високим рівнем ризику дефолту, характерного для економічного середовища сучасної України. За період з січня 2014 по січень 2017 роки кількість платоспроможних банків в Україні скоротилося майже в два рази – з 181 до 93. Це змушує господарюючих суб'єктів і фізичних осіб бути максимально обережними при виборі партнерів у фінансовій сфері [192].

Традиційний підхід до аналізу надійності контрагентів передбачає вивчення фінансових показників їх діяльності. Однак, коли кількість досліджуваних економічних систем зростає до сотень та тисяч, вирішити задачу таким методом за прийнятний термін стає практично неможливо.

Одним з існуючих непрямих методів її вирішення є використання інтегральних оцінок, одержуваних з початкових даних, шляхом обробки деякою математичною функцією. Вхідними даними у фінансовій сфері, як правило, є фінансова звітність, зокрема баланс та звіт про фінансові результати. Таким чином, можна отримати оцінки, що характеризують, наприклад,

прибутковість, надійність і стабільність кандидатів у партнери. Як за кордоном, так і в Україні створено чимало подібних систем. Наприклад, в США в 1970-х - 1990-х роках були створені та згодом отримали поширення у світовій практиці системи CAMEL, CAEL, CAMELS, FIMS, UBSS [118].

Світова фінансова криза, що почалася в 2008 році, стимулювала нові дослідження стійкості фінансово-кредитних систем. Так, розвитком ідеї створення інтегральних показників є роботи українських вчених Лернера [131], Боровского, Гатинского [93]. Однак, всім їм властивий ряд недоліків:

- жорстко заданий набір коефіцієнтів і «ваг» в моделях;
- кількість аналізованих показників і трудомісткість аналітичної роботи залишаються досить високими;
- утруднена оцінка розвитку об'єкта;
- порогові величини між ступенями рейтингу є усередненими і суб'єктивними.

Якщо сформулювати вимоги до способу аналізу надійності учасників ринку фінансових послуг, як антитезу до перелічених недоліків, то необхідний метод, легкий у використанні, наочний, гнучкий і самоналагоджувальний. Зрозуміло, повинна забезпечуватися достовірність отриманих результатів і оцінок.

Деякі з перелічених вимог реалізовані в методі кластерно-регресійного аналізу фінансових показників на основі групового обліку аргументів [208], проте легкість у використанні і адаптивність в цьому випадку недостатні.

Розглянемо варіанти застосування в якості інструменту аналізу надійності учасників ринку фінансових послуг (на прикладі банківської системи) методів нейромережевого моделювання.

Інформаційну базу дослідження становлять відкриті дані про активи, пасиви і капітал українських банків (всього 60 статей) станом на квітень 2014 року (тобто до початку масового банкрутства банків) і на жовтень 2015 року (після того, як пройшла перша хвиля банкрутств) [228]. Для дотримання порівнянності даних, відповідно до рекомендацій, даних в п. 2.1, замість абсолютних значень взято відносини відповідних показників до сумарної величини активів банку:

$$p_{i,j-1} = \frac{s_{i,j}}{s_{i,1}}, \quad i = \overline{1, n}; j = \overline{2, m}, \quad (3.3)$$

де n – кількість банків; m – кількість параметрів, що характеризують їх діяльність; s_{ij} – елемент матриці вхідних значень S ; p_{ij} – елемент матриці нормалізованих значень P . Перший стовпець матриці S – це сума активів банку.

Зазначимо, що такий спосіб подання даних не є єдиним. Наприклад, в [203] вхідну вибірку складають переважно розрахункові показники, що характеризують бізнес-моделі роботи банків.

Для моделювання рішення задачі використано наступні програмні продукти: Microsoft Excel (підготовка даних), Deductor Studio Academic

(моделювання і аналіз результатів).

Оскільки в даній задачі потрібно віднести досліджувані об'єкти до одного з двох класів – «банкрут», або «платоспроможний», очевидною базовою постановкою задачі є «бінарна класифікація». Розглянемо її рішення в цій постановці [149].

Процес моделювання представимо у вигляді наступних етапів:

Етап 3.1. Підготовка і попередня обробка даних.

На цьому етапі інформація перетворюється зі початкового вигляду в єдиний масив даних, який нормується за допомогою формули (3.3) і зберігається у форматі, необхідному для подальшої обробки програмою Deductor Studio. При підготовці даних відразу ж додається інформація про банки, діяльність яких протягом наступного часу була припинена (за даними [179, 230]). Слід зазначити, що в навчальній вибірці використовуються дані про показники роботи банків станом на квітень 2014 року та про банки, які стали банкрутами до жовтня 2015 року. Тестової вибіркою є показники роботи банків станом на жовтень 2015 року та інформація про банки, які стали банкрутами в період з листопада 2015 до березня 2017 року. Загальний обсяг такої вибірки становить 311 записів, з яких в навчальній перебуває 181 і в тестовій 130.

Етап 3.2. Відсіювання параметрів з низькою значимістю.

Навіть поверхневий аналіз вхідної вибірки даних, за формулами (2.5) – (2.7) показує недостатність її обсягу. Тому необхідно провести аналіз значимості вхідних показників та відсіяти ті з них, які слабо впливають на кінцевий результат. Для аналізу значимості можна використати методи, розглянуті в 2.1, а також аналітичні прийоми. За результатами аналізу виявилось, що такі статті, як «Незареєстровані внески до статутного капіталу», «Емісійні різниці», «Необоротні активи, утримувані для продажу» і деякі інші, вага яких в балансі є незначною, слабо впливають на платоспроможність банку. Для перевірки того, чи не було відсіяно значущі параметри, можна порівнювати ефективність роботи деякої ШНМ з фіксованою архітектурою при різних наборах вхідних даних.

Етап 3.3. Вибір архітектури ШНМ і аналіз результатів моделювання.

На цьому етапі слід розглянути різні архітектури ШНМ з точки зору якості навчання. Основним критерієм приймаємо середній рівень помилки на навчальній та тестовій множині (табл. 3.4).

Таблиця 3.4 – Результати навчання ШНМ різної архітектури в задачі прогнозування банкрутства банків України

№ _{пп.}	Архітектура ШНМ	Розпізнано на навчальній множині	Розпізнано на тестовій множині
1.	47-5-2	95,6%	72,3%
2.	47-3-2	96,7%	75,4%
3.	47-7-2	98,9%	76,15%
4.	47-5-3-2	96,1%	72,3%
5.	47-6-3-2	97,8%	76,92%

Аналіз табл. 3.4, показує, що результати ШНМ різної архітектури досить близькі. Тому для остаточного вибору скористаємося критерієм спряженості результатів (табл. 3.5-3.8).

Таблиця 3.5 – Таблиця спряженості результатів для ШНМ з архітектурою 47-3-2

	Класифіковано		
Фактично	Платоспроможний	Банкрут	Всього
Платоспроможний	85	9	94
Банкрут	23	13	36
Всього	108	22	130

Таблиця 3.6 – таблиця спряженості результатів для ШНМ з архітектурою 47-7-2

	Класифіковано		
Фактично	Платоспроможний	Банкрут	Всього
Платоспроможний	88	6	94
Банкрут	25	11	36
Всього	113	17	130

Таблиця 3.7

Таблиця спряженості результатів для ШНМ з архітектурою 47-5-3-2

	Класифіковано		
Фактично	Платоспроможний	Банкрут	Всього
Платоспроможний	78	16	94
Банкрут	20	16	36
Всього	98	32	130

Таблиця 3.8 – Таблиця спряженості результатів для ШНМ з архітектурою 47-6-3-2

	Класифіковано		
Фактично	Платоспроможний	Банкрут	Всього
Платоспроможний	88	6	94
Банкрут	24	12	36
Всього	98	32	130

Для даної задачі більшу вартість має помилка класифікації банків-банкрутів, як платоспроможних, ніж класифікація платоспроможних банків, як банкрутів. Спираючись на цей критерій, з аналізу таблиць спряженості (табл. 3.5-3.8) можна зробити висновок, що найгіршою з розглянутих є ШНМ з архітектурою 47-7-2, яка показала найкращі результати при навчанні. При перевірці на тестових даних ця мережа помилково класифікувала 25 банків-банкрутів, як платоспроможних. Кращою, за цим критерієм, виявилася ШНМ з

архітектурою 47-5-3-2, яка помилково віднесла до платоспроможних 20 банків-банкрутів з 36.

Таким чином результати застосування персептронної ШНМ для вирішення задачі прогнозування банкрутств комерційних банків в постановці бінарної класифікації можна вважати задовільними. Слід зазначити, що прогнозування здійснювалося на досить тривалий горизонт (майже 1.5 року від моменту зрізу даних). Крім того ситуація в банківській системі обумовлена не тільки фінансовими, а й політичними чинниками. Наприклад, майже всі розглянуті ШНМ при аналізі тестової вибірки класифікували Приватбанк, як потенційного банкрута (рис. 3.5).

№ з/п	Назва банку	Банкрут	COL4	Банкрут_OUT
172	БАНК ВЕЛЕС	☒	2015,9	☒
173	СХІДНО-ПРОМИСЛ. КОМЕРЦ. БАНК	☒	2015,9	☐
174	БАНК АВАНГАРД	☐		☐
175	БАНК ПОРТАЛ	☐		☐
176	"ЦЕНТР"	☐		☐
177	АЛЬПАРІ БАНК	☐		☐
178	ДЕРЖЗЕМБАНК	☐	2016	☐
179	"ГЕФЕСТ"	☐		☐
180	ВЕКТОР БАНК	☐	2017	☐
181	УКРАЇНСЬКИЙ БАНК РЕКОНСТР.ТА РО	☐		☐
1	ПРИВАТБАНК	☐	0	☒

Рисунок 3.5 – Витяг з таблиці з результатами класифікації

Подальший розвиток подій показав, що фінансовий стан Приватбанку справді не дозволяв продовжувати нормальну діяльність, і потрібна була масштабна докапіталізація банку державою. Однак в проаналізованих джерелах інформації про неплатоспроможні банки [179] Приватбанк відсутній. Там же, на рис. 3.5, видно дані по Східно-промисловий комерційний банк, який ШНМ «відмовилася» класифікувати як проблемний навіть на навчальній вибірці. Детальний аналіз показав, що причиною ліквідації цього банку є його знаходження на тимчасово-окупованих територіях в Луганській обл.

Проте, навіть з урахуванням додаткової інформації про чинники припинення діяльності банків, результати вирішення задачі прогнозування банкрутств, як задачі бінарної класифікації далекі від ідеальних. Розглянемо її рішення в іншій постановці – угруповання об'єктів в рамках задачі кластеризації.

Як і раніше, застосуємо для цього самоорганізаційні штучні нейронні мережі Кохонена. У порівнянні з іншими інтелектуальними методами аналізу даних, вони забезпечують такі вигоди:

- можливість роботи на відносно невеликому масиві даних;
- немає необхідності в експертних оцінках, що виключає суб'єктивність результатів;
- високий ступінь формалізації підготовки даних і навчання мережі;
- наочність інформації, виведеної у вигляді карт Кохонена.

В процесі моделювання Етап 3.1 і Етап 3.2 відповідають розглянутим вище, з невеликими поправками на особливості мереж Кохонена. Так, на Етапі 3.2. вибір показників може проводитися на підставі профілів кластерів, які містять інформацію про значущість кожного показника, що дозволяє відсіяти ті з них, які слабо впливають на кінцевий результат. В цьому сенсі Етап 3.2. частково перетинається з Етапом 3.3.

Етап 3.3. Вибір архітектури та навчання нейронної мережі. Вибір розмірності карти Кохонена здійснюється виходячи з обсягу вхідній вибірки та особливостей задачі. Оскільки інформація про банкрутство банків на вхід ШНМ Кохонена не подається, для навчання може використовуватися вся вибірка даних, яка в даному випадку містить 311 прикладів. Для відображення вибірки такого обсягу можна вибрати розмірність мережі 4x4 і кількість кластерів 4, хоча припустимі відхилення від цих значень. Далі необхідно провести підстроювання параметрів навчання для отримання мінімальних значень помилки. Результатом є візуальні кластерні карти вхідних даних.

Етап 3.4. Проекція на отримані карти множини даних, що характеризують банки – банкрути, аналіз результатів. Зупинимося на цьому етапі докладніше.

Для кожної комірки карти визначається, які банки з тих, що потрапили до нього, протягом 2014-2015 років припинили діяльність. Результати можуть бути представлені як у вигляді нової карти, так і у вигляді таблиці (табл. 3.9).

Таблиця 3.9 – Розподіл банків, що припинили діяльність, по коміркам карти

№ комірки	Кількість банків		частка банкрутів		№ комірки	Кількість банків		частка банкрутів
	всього	банкрутів				всього	банкрутів	
0	6	1	17%		8	2	2	100%
1	9	4	44%		9	7	0	0%
2	9	1	11%		10	19	10	53%
3	7	1	14%		11	28	10	36%
4	2	1	50%		12	19	7	37%
5	8	3	38%		13	26	15	58%
6	9	3	33%		14	10	2	20%
7	10	2	20%		15	9	0	0%

Аналіз табл. 3.9 показує, що як загальна кількість банків, так і кількість тих банків, що припинили роботу, в різних комірках карти суттєво відрізняється. Однак на карті можна виділити надійні, і проблемні зони. Розглянемо розподіл банків по коміркам карти Кохонена (рис. 3.6) і розташування на цій карті кластерів, які виділено алгоритмом кластеризації (рис. 3.7).

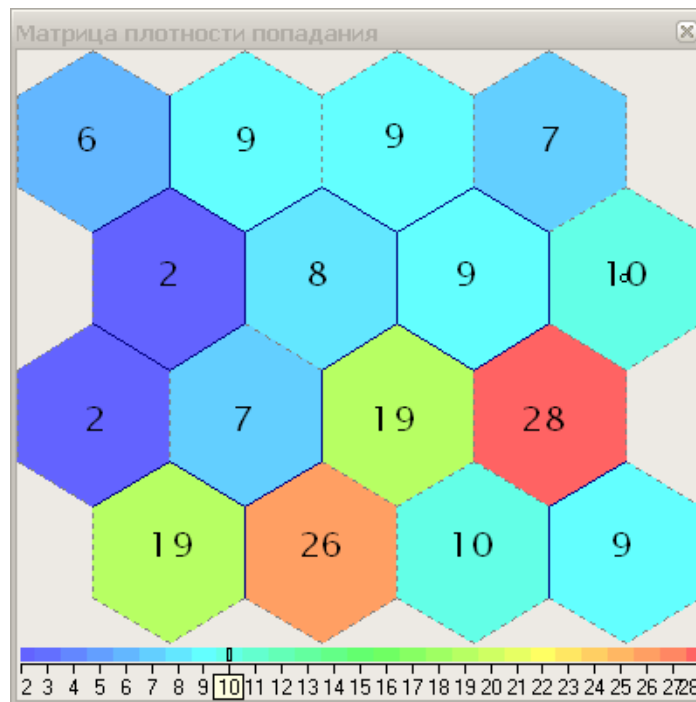


Рисунок 3.6 – Матриця щільності попадання банків в комірки карти Кохонена

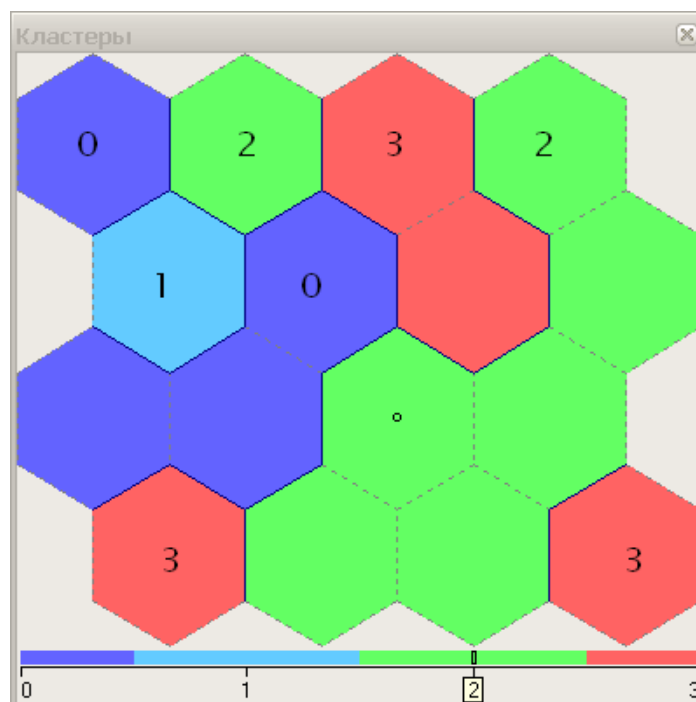


Рисунок 3.7 – Результати автоматичної кластеризації карти Кохонена

Нумерація комірок на рис. 3.6 та 3.7 здійснюється зліва направо та зверху вниз. Комірка з номером 0 розташовано у лівому верхньому куті.

Як можна помітити з сумісного аналізу рис. 3.6, рис. 3.7 та табл. 3.9, більшість банків, які припинили свою діяльність об'єднує кластер 2. Але за його межами також є комірки, які містять велику кількість банків – банкрутів. Отже, автоматична кластеризація не дозволяє знайти найкраще рішення даної задачі.

Ручна кластеризація передбачає покомірковий аналіз карти. Так, до достатньо надійних (за рівнем 20%) відносяться комірки 0, 2, 3, 7, 9, 14, 15. До проблемних – комірки 1, 4, 8, 10, 11, 12, 13.

З 58 банків, які потрапили в «надійні» комірки тільки 7 (12%) протягом 2014-2015 років припинили роботу. З 105 банків, що потрапили в «проблемні» комірки, 49 припинили працювати (47%).

Після проведення ручної кластеризації розглянемо результати обробки самоорганізаційною ШНМ на інформації про показники роботи банків, за станом на жовтень 2015 року. Результати розподілу банків по комірках карти показано в табл. 3.10.

Таблиця 3.10 – Розподіл банків, які припинили діяльність в 2016 р по комірках карти

№ комірки	Кількість банків		частка банкрутів		№ комірки	Кількість банків		частка банкрутів
	всього	банкрутів				всього	банкрутів	
0	2	2	100%		8	1	1	100%
1	9	4	44%		9	9	0	0%
2	2	1	50%		10	4	1	25%
3	6	2	33%		11	19	7	37%
4	1	1	100%		12	15	6	40%
5	6	1	17%		13	19	7	37%
6	4	1	25%		14	17	1	6%
7	9	1	11%		15	7	0	0%

Проаналізуємо розподіл банків з табл. 3.10 по кластерам, які виділено раніше.

Так в кластер «надійних» потрапляє 52 банки, з яких тільки 7 (13,5%) припинило свою діяльність в період з листопада 2015 до березня 2017 року.

У кластер «проблемних» потрапляє 68 банків, з яких 27 (40%) припинили свою діяльність в той же період.

Ще 10 банків попали в комірки 5 і 6, які за результатами попереднього аналізу неможливо віднести до жодної категорії.

Таким чином таблиця спряженості результатів прогнозування банкрутств з використанням самоорганізаційної мережі Кохонена має такий вигляд (табл. 3.11):

Таблиця 3.11 – Таблиця спряженості результатів для ШНМ Кохонена

	Класифіковано		
Фактично	Платоспроможний	Банкрут	Всього
Платоспроможний	45	41	86
Банкрут	7	27	34
Всього	52	68	120

З порівняльного аналізу табл. 3.5-3.8 та табл. 3.11 можна бачити, що застосування самоорганізаційних мереж дозволяє істотно знизити помилку першого роду, тобто класифікацію проблемних банків як надійних. Так, з 36 банків, що стали неплатоспроможними у розглянутий період, тільки 7 було раніше класифіковано, як «хороші». Саме ця помилка має найбільшу економічну значущість в даній задачі.

Важливі результати дає також аналіз окремих комірок карти. Так, на підставі табл. 3.9 і 3.10, найбільш надійними слід вважати банки, що потрапили в комірки 9 і 15, так як за весь аналізований період жоден з них не припинив існування. Причому, якщо в комірці 9 розташовуються банки, які переважно належать до групи малих, то в комірці 15 нейронна мережа відносить банки з категорії лідерів галузі. На початок 2016 року серед них були такі банки, як Промінвестбанк, Сбербанк, ВТБ, ІНГ, Піреус, Віес і Кредит-Європа. Аналогічно може бути виконано аналіз і за проблемними комірками.

Таким чином, в задачі прогнозування банкрутств комерційних банків більш вдалою є її постановка, як задачі угруповання об'єктів в рамках кластеризації. В рамках цієї постановки використання самоорганізаційних нейронних мереж дозволяє отримати достатньо достовірні оцінки, які можна використовувати в практичній діяльності для визначення надійності потенційних фінансових партнерів. Хоча отримання абсолютно достовірного прогнозу в поточних економічних і політичних умовах не є можливим, запропонований метод може грати роль одного з індикаторів фінансової стійкості при виборі банків, страхових компаній та інших фінансових посередників.

В цілому, отримані результати підтверджують Твердження 2.1 про те, що ефективність вирішення складних економічних задач залежить від методів і інструментів інтелектуальних обчислень, які застосовуються для цього.

3.2 Генетичні моделі обробки даних в економічних системах

Дослідження динаміки появи і накопичення інформації в сучасному суспільстві об'єктивно показують, що воно знаходиться в стані, яке отримало назву «Інформаційний вибух». Масове впровадження комп'ютерних технологій призвело до того, що кількість інформації, яка потребує обробки, збільшується в геометричній прогресії приблизно на 30% щорічно [25, с. 63]. У цих умовах збільшується актуальність методів, що дозволяють спростити сприйняття даних, а також їх зберігання та передачу.

Проблема спрощення даних, представлених у вигляді динамічних рядів, до сих пір досліджувалась в основному стосовно теорії передачі інформації і цифровій обробці аналогових сигналів. Крім того, схожі методи існують в рамках економічної статистики і в деяких інших дисциплінах. Однак всі вони розроблялися в умовах дефіциту обчислювальних ресурсів і тому реалізують лише прості прийоми, не дозволяючи одночасно отримати більшу ступінь стиснення даних і зберегти їх адекватність [76].

Можливості сучасних інформаційних технологій дозволяють розробити ефективний *метод спрощення динамічних рядів*. Далі розглянуто вирішення цієї проблеми, яке включає [155]:

- аналіз переваг і недоліків існуючих методів;
- формування принципів, що дозволяють уникнути цих недоліків;
- побудову математичної моделі пошуку оптимального представлення даних з різним ступенем деталізації;
- практичну реалізацію даної моделі з використанням генетичних алгоритмів в системі Matlab;
- аналіз результатів моделювання.

Слід відмітити, що в сучасному суспільстві темпи зростання обсягів інформації випереджають темпи розвитку засобів її обробки. Внаслідок цього на конкурентоспроможність економічних суб'єктів і соціальних груп істотно впливає як нерівність доступу до засобів інформаційних технологій, так і нерівність в знаннях про використання таких технологій [50 с. 5]. Саме з проблемою обробки великих обсягів інформації пов'язаний бурхливий розвиток технологій обробки великих обсягів даних (*Big Data*), які дозволяють знаходити закономірності в базах з мільйонами записів. Тому черговий пік популярності переживає такий інструмент інтелектуальної обробки інформації, як нейронні мережі, можливості яких по розпізнаванню графічних і звукових об'єктів в даний час досягли людських, а в деяких випадках навіть перевершують їх. Тут, по суті, основною задачею, яка вирішується, є зниження розмірності даних до рівня, достатнього, для подальшої обробки традиційними методами аналізу, або для безпосереднього сприйняття ОПР. Те ж призначення має і спрощення динамічних рядів, що показують зміни в послідовні періоди часу розмірів будь-якого явища або ознаки [136, с.85, с.53]. Залежно від складових величин розрізняють кілька типів динамічних рядів:

Базовими є динамічні ряди, побудовані з абсолютних величин (курс валюти, обсяги реалізації товару і так далі). Решта типів виходять шляхом обробки абсолютних показників, а саме:

- *похідні динамічні ряди*, представлені відносними величинами і демонструють зміни будь-яких коефіцієнтів (зміна цін, зміна обсягів реалізації і так далі);
- *динамічні ряди, що складаються з середніх величин*, наприклад, показників середньої реалізації товарів за період часу, середнього доходу на душу населення і так далі.

Кожен динамічний ряд складається з двох елементів: відрізка часу (періоду), в рамках якого було зафіксовано певний показник і показника, що характеризує об'єкт дослідження (рівні ряду). Якщо стан об'єкта описується m показниками, то масив даних A буде мати вигляд:

$$A = \{t_i, a_{i1}, a_{i2}, \dots, a_{ij}\}, \quad i = 1..m, j = 1..n$$

У деяких випадках, якщо відлік показників ведеться з однаковими

інтервалами часу, елемент t_i в масиві даних може бути пропущений, оскільки легко відновлюється з t_0 , i та Δt , які зазвичай відомі.

Сучасні інформаційні технології дозволяють отримувати дані з високою частотою, що дозволяє краще відстежувати зміни в стані об'єктів, але ускладнює сприйняття отриманих даних та їх подальшу обробку, так як вони містять велику кількість надлишкової інформації – шуму. Такі дані можуть бути оброблені з метою видалення невеликих змін показників і виділення значних тенденцій. Оскільки в результаті кількість точок відліку (m) зменшується, можна говорити про *спрощення* даних.

Серед існуючих методів вирішення задачі спрощення даних можна виділити як загальні, так і спеціальні. До загальних відноситься, наприклад, метод квантування за рівнем – розбиття діапазону значень неперервної або дискретної величини на кінцеве число інтервалів. (рис. 3.8).

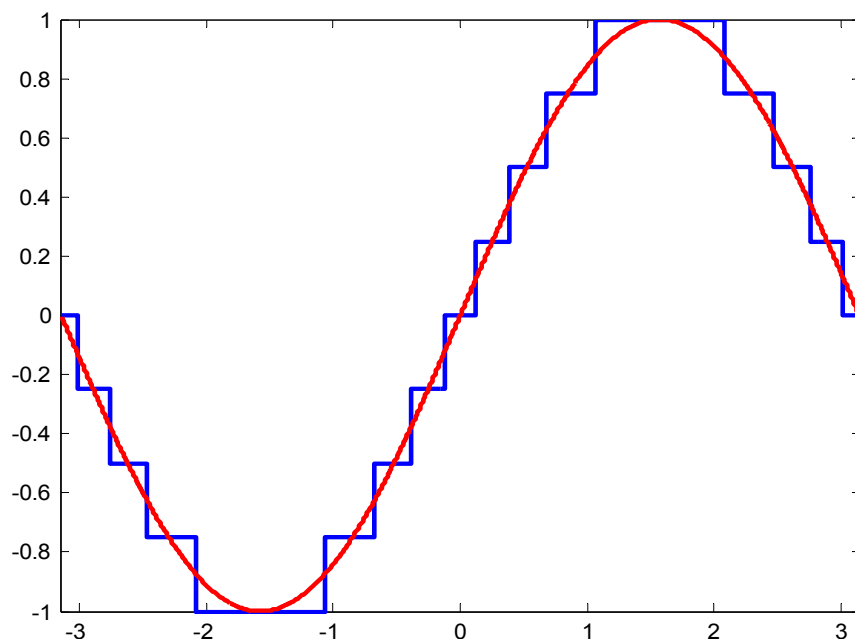


Рисунок 3.8 – Спрощення динамічного ряду методом квантування за рівнем

Як можна побачити з аналізу рис. 3.8, якщо протягом деякого періоду часу значення величин динамічного ряду не виходило за межі одного інтервалу, то в результаті квантування всі ці величини будуть замінені одним значенням [109, с. 184].

Перевагою алгоритму квантування за рівнем є простота реалізації і мінімальні вимоги до обчислювальних ресурсів, що зумовило його поширення в методах цифрової обробки сигналів. Недолік методу – сильні спотворення даних при збільшенні стиснення.

Спеціальні методи застосовуються у вузьких предметних областях і максимально враховують їх особливості. Так, для стислого представлення біржових даних використовується метод квантування за часом, модифікований так, що кожен дискретний часовий інтервал описується не одним, а чотирма значеннями (ціни відкриття і закриття, а також максимальна і мінімальна ціна

за період).

Приклад такого графіка наведено на рис. 3.9.

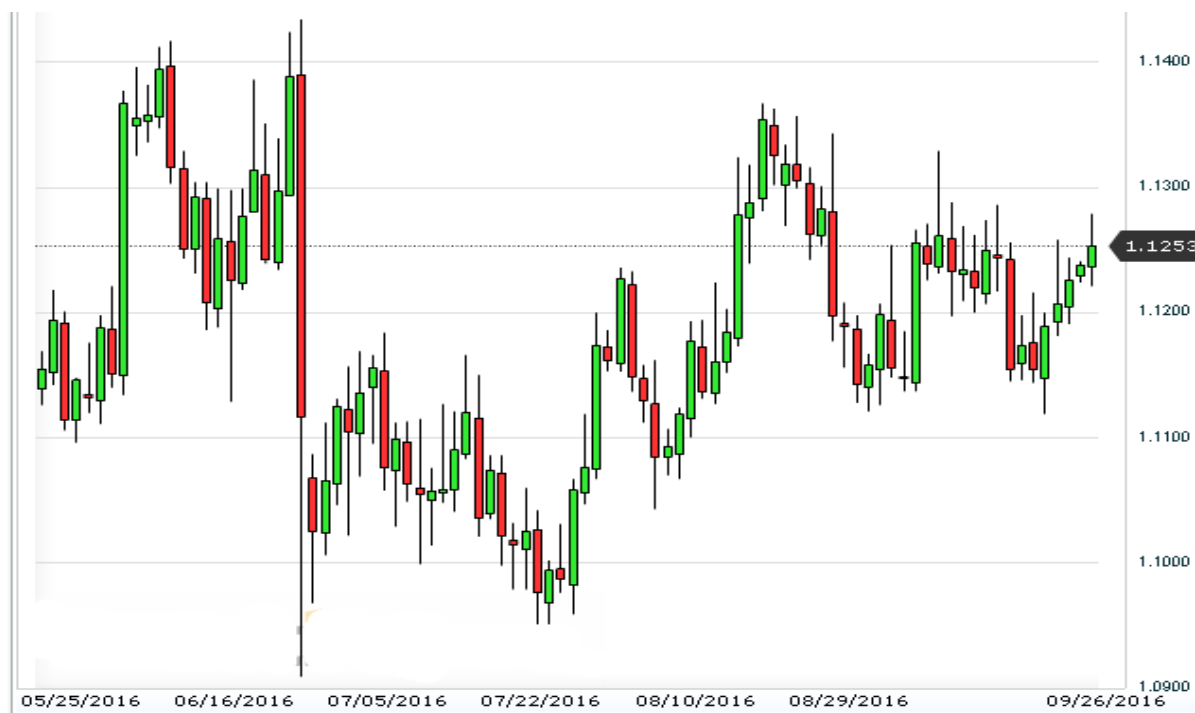


Рисунок 3.9 – Ілюстрація принципів квантування за часом з постійним кроком

Для біржових даних існує кілька стандартних часових інтервалів, наприклад – 1 година, 4 години, 1 день. Цей метод зручний тим, що незалежно від ступеня стиснення, зберігаються найважливіші, з точки зору біржових операцій, характеристики часового ряду. Крім того, результати обробки можуть бути наочно представлені у вигляді "японських свічок" (рис 3.9), або графіка із зарубками [115].

Разом з тим, даний метод є вузькоспеціалізованим, оскільки розрахований на обробку тільки одного показника – ціни, а крім того вимагає від людини певної підготовки для сприйняття результатів.

Загальним недоліком розглянутих методів є наявність фіксованого кроку квантування, що не дозволяє отримати високу ступінь стиснення даних при збереженні достатнього рівня апроксимації.

Також можна відзначити ряд методів механічного згладжування динамічних рядів, привнесених зі статистики. Наприклад, метод укрупнення інтервалів, аналітичного вирівнювання, згладжування рядів за допомогою ковзних середніх та подібні. Всі вони також не позбавлені недоліків, основним з яких є згладжування та втрата важливих для аналізу пікових рівнів, а також втрата деяких закономірностей в даних [84]. Крім того, в результаті згладжування не завжди відбувається скорочення точок відліку значень ряду.

Більшості з зазначених недоліків дозволяє уникнути метод *квантування за часом зі змінним кроком*, який пропонується далі.

У загальному вигляді завдання, що вирішується можна сформулювати наступним чином:

Визначення 3.1

Складністю динамічного ряду назвемо кількість точок відліку m , які необхідні для його формування.

Визначення 3.2

Спрощенням динамічного ряду назвемо таке перетворення

$$A \rightarrow \bar{A}, \quad (3.4)$$

при якому

$$\bar{m} < m, \quad (3.5)$$

а мера відповідності рядів A і $\bar{A}$,

$$\Psi(A, \bar{A}) \rightarrow \max. \quad (3.6)$$

де $A = \{t_i, a_i\}$, $i = 1..m$ – масив, який задає послідовність даних, що утворюють початковий динамічний ряд; $\bar{A} = \{t_i, a_i\}$, $i = 1..\bar{m}$ – результуючий масив даних; m та $\bar{m}$ – складність рядів.

Однак для практичного використання умов (3.4) – (3.6) недостатньо, оскільки очевидним рішенням буде ряд, що відрізняється від початкового тільки на один елемент. Тому необхідно ввести додаткову умову мінімізації точок відліку:

$$\Theta(\bar{m}) \rightarrow \min, \quad (3.7)$$

де Θ – міра складності ряду.

Вирази (3.4) - (3.7) складають закінчену модель спрощення ряду, але спосіб визначення функцій Ψ і Θ вимагає додаткових пояснень.

Функція Ψ має економічний сенс премії за відповідність рядів A і $\bar{A}$. При цьому відповідність виражається в збігу трендів в рядах і залежить від сили тренда.

Функція Θ має економічний сенс штрафу за кожну точку відліку. Варіюючи розмір штрафу можна впливати на відносну складність ряду $\bar{A}$.

Розглянемо практичні приклади таких функцій.

Нехай A – динамічний ряд, що містить m пар елементів $\{t_i, a_i\}$, причому

$$\begin{aligned} a_i &\in [0;1], \\ t_i &= i. \end{aligned} \quad (3.8)$$

Якщо значення елементів a_i не відповідають умовам (3.8), то необхідно

провести нормування:

$$\bar{a}_i = \frac{a_i - \min(A)}{\max(A) - \min(A)}, \quad (3.9)$$

а якщо дані в ряді розташовуються не через рівні проміжки часу, то інтерполяцію, де проміжні елементи ряду між елементами a_i та a_{i+1} можуть бути знайдені за формулами лінійної інтерполяції:

$$\frac{t - t_i}{t_{i+1} - t_i} = \frac{a - a_i}{a_{i+1} - a_i}, \quad (3.10)$$

$$a = a_i + \frac{t - t_i}{t_{i+1} - t_i} (a_{i+1} - a_i), \quad (3.11)$$

або за формулами квадратичної, чи поліноміальної інтерполяції, які є більш складними, але дозволяють отримати більш гладку поверхню функції.

Оскільки рішення задачі зводиться до визначення точок перегину ряду, введемо в модель вектор рішень системи

$$S = \{s_j\}, \quad j = 1.. \bar{m},$$

де s_j – індексний елемент, чисельно рівний елементу t_i початкового ряду A , тобто місцю розташування точки перегину.

Тоді функція відповідності Ψ тоді запишеться наступним чином:

$$\Psi = \sum_{j=2}^{\bar{m}} |a_{t_j} - a_{t_{j-1}}|, \quad (3.12)$$

$$t_1 = 1,$$

а функцію штрафів можна записати так:

$$\Theta = l \cdot \bar{m}, \quad (3.13)$$

де l – штраф за кожну точку відліку.

У загальному вигляді цільова функція моделі має вигляд:

$$\max z = \Psi - \Theta. \quad (3.14)$$

Очевидно, що найбільшою відповідністю в загальному випадку має динамічний ряд, який відрховується в усіх тих же точках, що і початковий.

Але з іншого боку для нього найбільшою буде і функція штрафів. При цьому чим більше різниця між значеннями результуючого ряду в сусідніх точках відліку, тим більше значення цільової функції. Таким чином, рішення є компромісом між відповідністю із вихідними даними та кількістю точок відліку з іншого боку.

Процес рішення задачі зводиться до знаходження оптимального розміру і значень вектора S . Єдиним шляхом знаходження глобального оптимуму в просторі рішень даної задачі є повний перебір всіх можливих варіантів. Оскільки обчислювальна складність такої задачі зростає експоненціально із збільшенням кількості показників m , то вона відноситься до класу NP-важких і тому може бути вирішена лише із використанням евристичних методів багатфакторної оптимізації, які дозволяють за прийнятний час знайти приємне рішення. До них відносяться методи, що вже розглядалися: генетичні алгоритми, метод імітації відпалювання, метод мурашиних колоній та ін. В даному випадку розглянемо застосування генетичних алгоритмів.

Модель (3.12) – (3.14) реалізується в системі Matlab з використанням функцій пакета розширень Genetic Algorithm and Direct Search Toolbox. Вхідними параметрами програми є сам динамічний ряд, значення штрафу за точку відліку і деякі параметри самого генетичного алгоритму (розмір популяції, умови зупинки еволюції, ймовірність мутацій).

Результати моделювання представляються як в графічному вигляді, так і у вигляді значення функції пристосованості кращої особи в популяції. Зважаючи на особливості даного пакета математичного моделювання, задачі оптимізації в ньому можуть вирішуватися тільки на мінімум. Тому при адаптації моделі до програмного середовища Matlab цільову функцію визначимо як:

$$\min z = \Theta - \Psi, \quad (3.15)$$

а також зробимо деякі інші не принципові спрощення.

В якості вхідних даних, на яких тестувався алгоритм, обрано дані про зміну курсу євро по відношенню до долара США за січень 2016 роки (120 точок відліку) і дані щодо зміни денної температури в м. Маріуполь з січня по вересень 2016 роки (259 точок відліку). Для полегшення візуального аналізу якості роботи алгоритму, на графіках присутні як початковий, так і результуючий ряди.

На рис. 3.10 і рис. 3.11 показано графіки зміни показника денної температури, оброблені алгоритмом з різним ступенем стиснення.

Як бачимо, навіть при невеликому ступені стиснення, при збереженні практично всіх основних тенденцій, які спостерігаються на графіку, кількість точок відліку знижується з 259 до 30, тобто більше ніж у 8 разів.

Подальше підвищення рівня стиснення дозволяє знизити кількість точок відліку до 16 (рис. 3.11). При ще більшому стисненні кількість точок відліку доходить до 3 (початок ряду, точка максимуму і кінець ряду). Але навіть в цьому випадку даних достатньо для відстеження основних тенденцій зміни

вхідного показника.

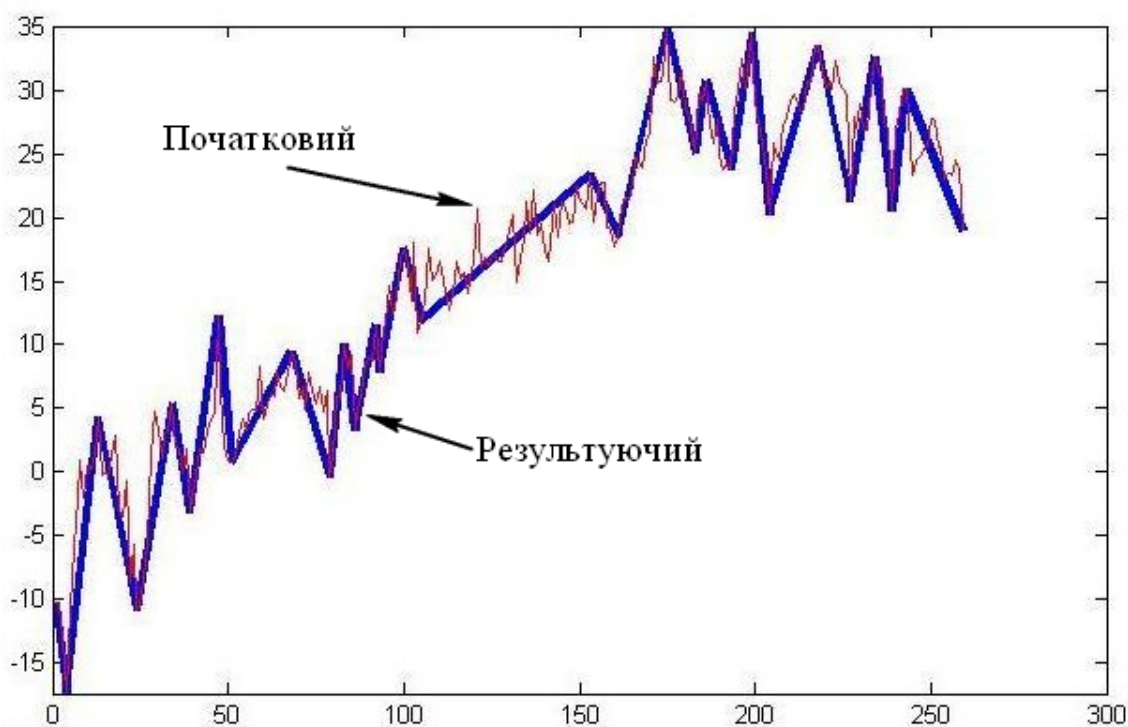


Рисунок 3.10 – Початковий і спрощений графіки зміни показника денної температури в січні-вересні 2016 р при ступені стиснення $l=0.1$

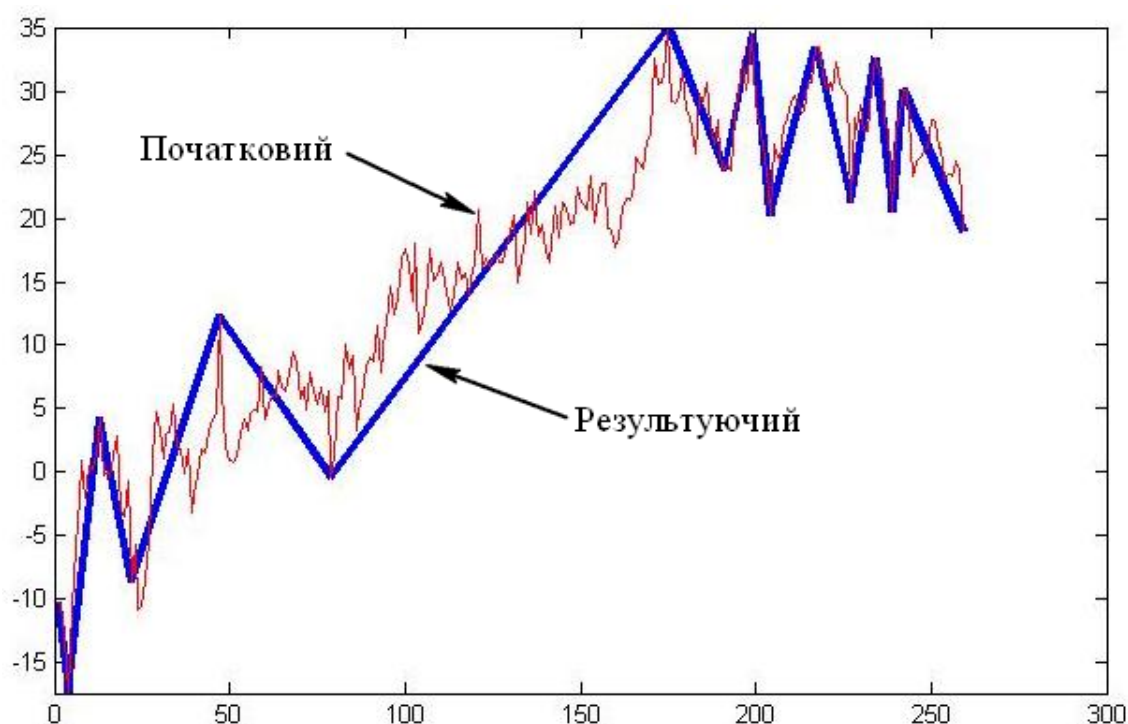


Рисунок 3.11 – Початковий і спрощений графіки зміни показника денної температури в січні-вересні 2016 р при ступені стиснення $l=0.2$

Розглянемо результати, які було отримано при обробці даних про коливання валютного курсу (рис. 3.12).

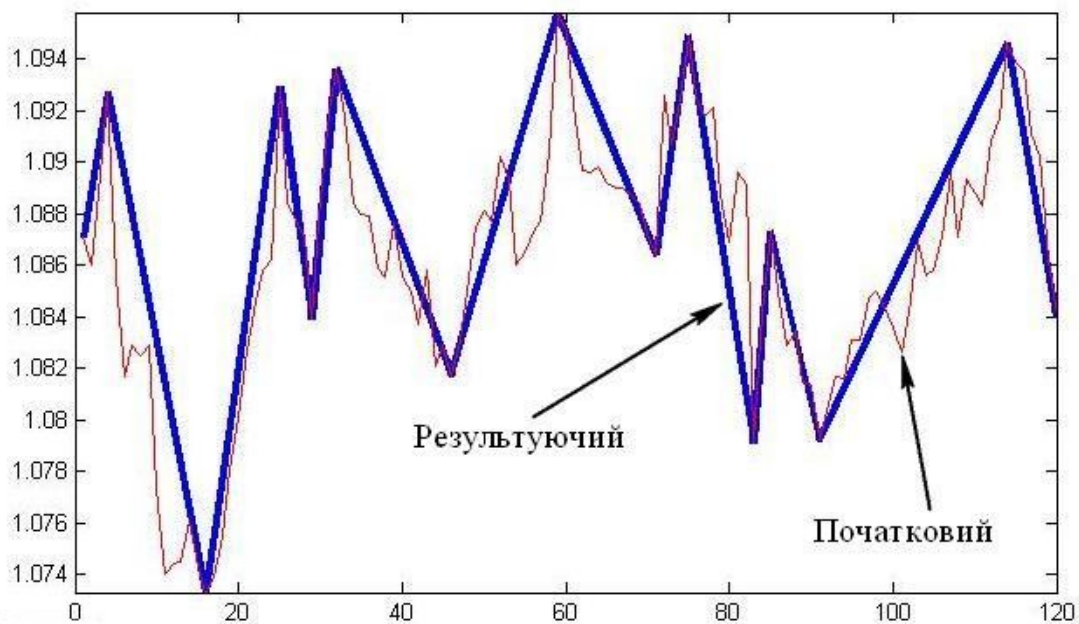


Рисунок 3.12 – Початковий і спрощений графік коливань курсу EUR / USD в січні 2016 р при ступені стиснення $l=0.2$ и рівні пристосованості 4.5467

При $l=0.2$ генетичним алгоритмом було виділено основні точки перегину, загальна кількість яких скоротилася з 120 до 15 (в 8 разів). При цьому час роботи програми на комп'ютері із процесором Athlon 64 4400+ склав близько 5 секунд.

Слід зазначити, що рішення, які знайдені при різних запусках алгоритму, можуть дещо відрізнятися одне від одного і відповідно мати різні значення функції пристосованості. Для перевірки того, наскільки сильно розрізняються знайдені рішення, було зроблено 20 тестових запусків генетичного алгоритму по обробці коливань валютного курсу. Гістограма розподілу значень функції пристосованості показана на рис. 3.13.

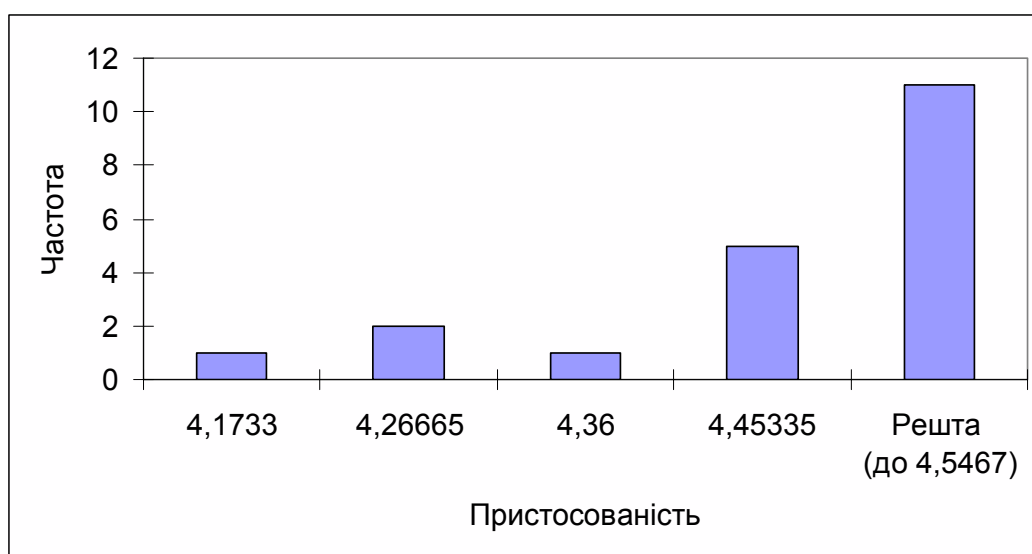


Рисунок 3.13 – Розподіл значень функції пристосованості при 20 запусках генетичного алгоритму

З аналізу гістограм (рис. 3.13), можна зробити наступний висновок: В 11 випадках з 20 значення функції пристосованості практично не відрізнялися від кращого, ще в 5 випадках значення функції пристосованості відрізнялося від одного з найкращих знайдених не більше ніж на 0.1. І тільки в 4 випадках були зафіксовані більш значні відхилення значень функції пристосованості. Причому додатковий аналіз показав, що навіть гірше зі знайдених рішень адекватно відображає основні тенденції ряду і містить тільки дві невеликих помилки у визначенні позиції точок перегину (рис. 3.14).

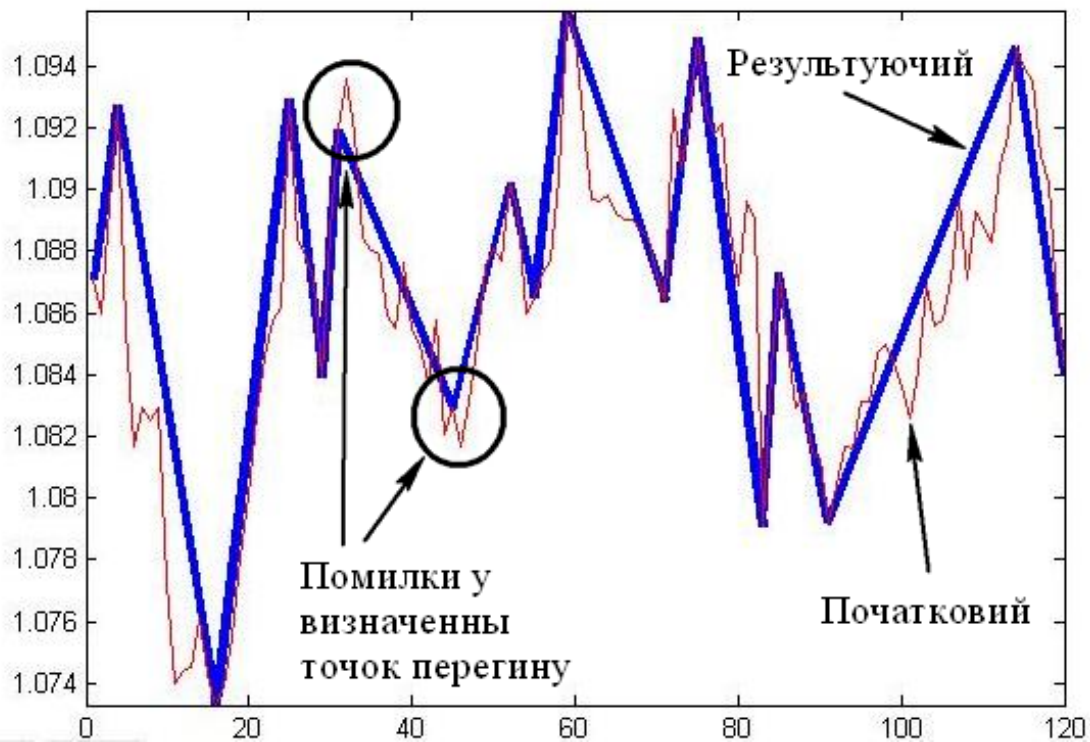


Рисунок 3.14 – Початковий і спрощений графік коливань курсу EUR / USD в січні 2016 р при ступені стиснення $l=0.2$ і рівні пристосованості 4.1733

Таким чином, запропонований метод спрощення динамічних рядів дозволяє швидко і ефективно скорочувати кількість точок відліку ряду з будь-яким ступенем стиснення даних, який регулюється зміною лише одного показника. У порівнянні з існуючими методами квантування він забезпечує збереження пікових значень ряду, а також більш ефективне стиснення даних, особливо якщо значення показників ряду змінюються повільно.

Метод спрощення динамічних рядів з використанням генетичних алгоритмів може використовуватися в економічній статистиці і аналізі, в процесі підготовки презентацій, а також у всіх інших випадках, коли потрібно виділити основні тенденції у великій послідовності даних. Окремо слід відзначити можливість застосування принципу квантування за часом зі змінним кроком в системах передачі інформації з обмеженою пропускнуою спроможністю.

3.3 Оцінка параметрів економічних систем методами системно-динамічного імітаційного моделювання

Особливості системно-динамічного моделювання, як одного з методів інтелектуальних обчислень було розглянуто в 1.2. Моделювання системної динаміки не тільки допомагає відповісти на питання про майбутній стан об'єкта, але і дозволяє краще зрозуміти його внутрішню структуру і взаємозв'язки між елементами. Системна динаміка передбачає найвищий рівень агрегування компонентів з усіх методів імітаційного моделювання. Завдяки низці спрощень, прийнятих в моделях (абстрагування від індивідуальних характеристик об'єктів і фізичних характеристик навколишнього середовища, безперервність всіх змінних і процесів), вони дозволяють простими засобами отримати адекватний опис процесів в досить складних системах.

У банківській діяльності одним з напрямків застосування системної динаміки є моделювання ризику в активних операціях. Це обумовлено тим, що хоча самі фактори ризику можуть бути виявлені емпірично, аналіз сукупного впливу на досліджувану систему і його змін у часі являє собою вкрай складну задачу, рішення якої традиційними методами практично неможливо.

Розглянемо приклади використання методів системної динаміки для моделювання зовнішніх і внутрішніх факторів ризику в банківських активних операціях.

Моделювання зовнішніх факторів ризику банківських активних операцій, на прикладі ціноутворення ринку житлової нерухомості. Згідно з оцінками [3], на початок 2016 року на нерухомість припадало близько 60% світових традиційних активів. Їх сукупна вартість (\$ 217 трлн.) в 2.7 рази перевищувала світовий ВВП. При цьому 75% загальної вартості нерухомості, або \$ 162 трлн. припадає на житло.

Довгий час інвестиції в ринок нерухомості вважалися одним з найнадійніших способів розміщення коштів. Невисока прибутковість операцій з нерухомістю компенсувалася їх низьким ризиком. Слабка волатильність ринку обумовлювала можливість застосування для його аналізу простих статистичних моделей, а дослідження ринку нерухомості часто були спрямовані на уточнення коефіцієнтів таких моделей.

У ХХІ сторіччі ситуація різко змінилася. Різне зростання цін на нерухомість змінилося різким обвалом. Обидві події не могли бути враховані за допомогою традиційних для цього ринку методів прогнозування. Багато банків, які використовували традиційні підходи для оцінки ризику операцій з нерухомістю, збанкрутували.

Світова фінансова криза, що почалася саме з ринку нерухомості, призвела до використання нових методів та інструментів його аналізу і прогнозування. Одним з них є моделювання системної динаміки. Використання цього методу для аналізу міської забудови, було запропоновано ще в 1969 році його автором Дж. Форрестером [19]. На початку 1970-х років ці ідеї спробували розвинути, але до практично значущих моделей так і не дійшли, оскільки ринок нерухомості тоді був стабільним і традиційні підходи описували його краще.

З початком світової фінансової кризи кількість публікацій про застосування імітаційного моделювання для аналізу ринку нерухомості зростає. Слід зазначити роботи корейських вчених з моделювання національного ринку нерухомості [9, 31], що відрізняються глибокою обробкою і великою кількістю врахованих факторів. Проміжним підсумком цього періоду можна назвати роботу [16], що вийшла в 2014 році, де систематизовано основні результати використання методів системної динаміки в дослідженні ринку нерухомості.

У вітчизняній літературі для дослідження ринку нерухомості досі використовуються традиційні для економетрії методи і моделі – факторні, статистичні та подібні їм [7, 94]. Однак, слід зазначити, що в Україні ускладнено доступ до багатьох соціально-економічних показників розвитку суспільства та інфраструктури, що перешкоджає практичній перевірці результатів наукових досліджень. Зокрема, через відсутність прямих даних для отримання оцінки доходів різних груп населення доводиться використовувати непрямий метод, заснований на гіпотезі логнормального розподілу доходів у суспільстві [154].

Розглянемо існуючі теоретичні підходи до ціноутворення на ринку житлової нерухомості.

Одним з них є синтетичний метод, при якому ціна розглядається, як інтегральна оцінка вартості різних факторів. Наприклад, ринкова ціна комерційної нерухомості може розглядатися, як похідна від інвестиційної, спеціальної, ліквідаційної, податкової та інших цін. Для житлової нерухомості цими факторами можуть бути попит, корисність, дефіцитність, можливість відчуження [160].

Інший підхід передбачає більшу кількість факторів, розподілених на три ієрархічних рівня: державний (соціальні, економічні, політичні), локальний (особливості місця розташування об'єкта, умови продажу) та індивідуальний (архітектурно-будівельні та фінансово-експлуатаційні).

Зважаючи на складність визначення перелічених факторів, на практиці для оцінки вартості нерухомості часто використовуються інші підходи, зокрема – витратний, дохідний, порівняльний. Останній з них найбільш поширений на ринку вторинного житла і, по суті, означає, що орієнтиром для встановлення нових цін будуть ціни, за якими аналогічне житло продавалося раніше. З огляду на те, що ціна завжди знаходиться в деякому «коридорі», ця модель може пояснити як зростання, так і падіння цін на ринку. При зростанні продавці орієнтуються на верхній кордон «коридору цін», а при падінні, відповідно на нижній. Встановлено, що в короткостроковому періоді еластичність цін на ринку нерухомості низька, а в довгостроковому – навпаки. Це означає, що ціни на ринку змінюються повільно, але в дуже широких межах, а процес зміни ціни є наслідком дії класичних законів попиту і пропозиції. Зазначені особливості ціноутворення на ринках житлової нерухомості покладено в основу моделей, які розглянуто далі.

В якості безпосереднього об'єкту для перевірки результатів моделювання обраний ринок вторинної нерухомості у м. Київ. Цей вибір обумовлено тим, що даний ринок є найбільшим в Україні (в житловому фонді Києва більш ніж 1

млн. квартир), а також наявністю доступної статистики цін на нерухомість за значний часовий інтервал [217].

Інформаційну базу дослідження складають дані державної служби статистики України (інформація про середні доходи населення), статистика цін на квартири в Києві, дані комерційних банків щодо умов кредитування на нерухомість, дані аналітичних агентств про обсяги операцій на ринку нерухомості. Для попередньої обробки даних використовуються методи математичної статистики [7].

Розглянемо системно-динамічну модель, яка реалізує процеси досягнення балансу між попитом і пропозицією. Аналіз статистики продажу житлової нерухомості в м. Київ показало, що пропозиція на вторинному ринку є практично постійною. Це дозволило зробити висновок, що ціна на ринку визначається переважно платоспроможним попитом і істотно спростити модель причинно-наслідкових зв'язків, яку наведено на рис. 3.15.

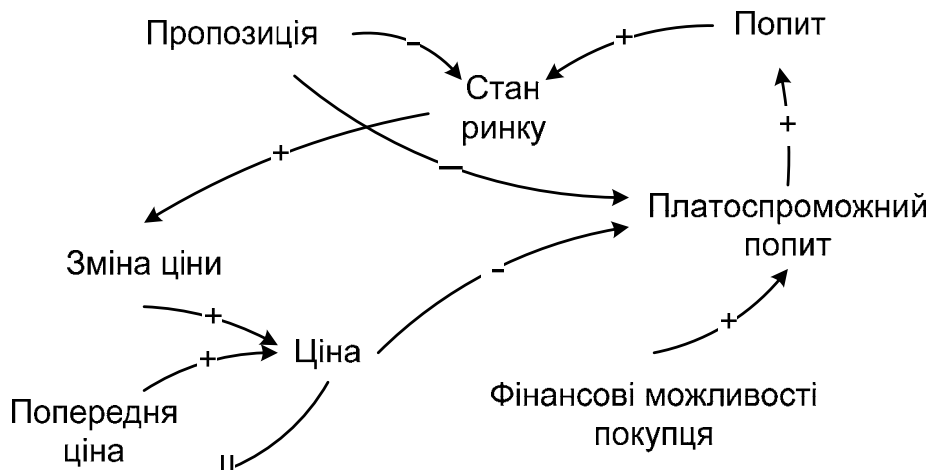


Рисунок 3.15 – Модель причинно-наслідкових зв'язків на ринку житлової нерухомості

Розглянемо цю модель. Як вже зазначалося, пропозиція в моделі є зовнішнім фактором і задається до початку імітаційного експерименту. При цьому пропозиція не обов'язково повинна залишатися постійною протягом всього системного часу, але всі її зміни повинні бути запрограмовані заздалегідь. Вважається, що пропозицію виражено в матеріальних одиницях (штуках).

Попит в моделі є змінною величиною, яку безпосередньо визначає платоспроможний попит. Чисельно ці показники тотожні. В моделі їх розділено для зручності подальшого розширення. Як і пропозиція, попит та платоспроможний попит є матеріальними величинами.

Відношення попиту до пропозиції визначає стан ринку і орієнтири для встановлення нових цінових рівнів. При цьому реалізовано порівняльний підхід, тобто орієнтиром для встановлення нових цін є ціни, за якими квартири продавалися раніше, які розташовані в деякому ціновому інтервалі – «коридорі». При визначенні параметрів цього коридору за основу обрано

нормальний закон розподілу ймовірностей, при якому найбільша кількість квартир продається за ціною, близькою до середньої. При невеликих змінах балансу попиту і пропозиції, коливання цін також будуть незначними, і тільки великі зміни викликатимуть коливання цін в межах всього коридору.

Ціна квартири в моделі визначається виходячи з попереднього значення ціни та її змін, викликаних змінами ринкового стану. Поточна ціна на наступному модельному кроці через лінію затримки стає попередньою.

Показник фінансових можливостей покупців відображає суму, яку потенційний покупець в середньому готовий заплатити за квартиру. Як і ціна, цей показник фігурує в моделі в грошовому вираженні. Численні дослідження ринку нерухомості свідчать, що реальними покупцями на ринку нерухомості є забезпечені люди, які складають лише 5-10% всього населення. Моделювання даного показника буде розглянуто нижче.

Якщо фінансові можливості покупців вище середнього рівня цін, це веде до зростання попиту, ринок переходить в стан дефіциту і починається поступове підвищення цін, яке триватиме до моменту їх вирівнювання з фінансовими можливостями покупців. В іншому випадку відбудеться зворотній процес [160].

Модель, яка подано на рис. 3.15, здатна адекватно описувати поведінку ринку нерухомості, але в силу самого визначення моделі, вона містить деякі спрощення і обмеження:

1. *Стабільна пропозиція.* Як уже зазначалося, пропозиція квартир на ринку вважається в моделі зовнішнім фактором і не залежить від стану ринку, цін та інших факторів. Хоча в реальності це не зовсім так, але гіпотеза про низьку еластичність пропозиції на ринку нерухомості дає підстави прийняти подібне припущення.

2. *Гомогенність пропозиції.* У моделі всі квартири мають схожі характеристики і відрізняються тільки ціною, в рамках закону нормального розподілу. У реальності пропозиція нерухомості досить різноманітна і покупці, в яких не вистачає грошей на велику квартиру в престижному районі, можуть придбати квартиру меншої площини, або вибрати менш престижне місце. Усунення цього обмеження суттєво ускладнить модель, тому слід обмежити область її застосування аналізом даних, де умова гомогенності виконується, наприклад однокімнатні квартири в певному районі.

Як можна помітити, основним зовнішнім фактором, що впливає на зміни показників в моделі, показаної на рис. 3.15, є фінансові можливості покупців. Розглянемо процедуру їх визначення.

Платоспроможний попит на ринку нерухомості формується в основному на підставі таких чинників, як власні кошти населення і залучені кошти. Власні кошти утворюються в результаті накопичення грошей, які населення отримує у вигляді заробітної плати та інших доходів. В якості залучених коштів будемо розглядати банківські іпотечні кредити, за рахунок яких під час буму операцій з нерухомістю купувалося до 80% квартир [112].

Модель причинно-наслідкових взаємозв'язків при визначенні фінансових можливостей покупців приведена на рис. 3.16.



Рисунок 3.16 – Модель причинно-наслідкових зв'язків формування фінансових можливостей покупців на ринку нерухомості

Базою формування фінансових можливостей, як можна простежити на рис. 3.16, є *доходи населення*. У багатьох країнах світу статистика грошових доходів населення ведеться розширено – з розбивкою за рівнями доходу. В Україні Державна служба статистики до 2016 року давала лише усереднене значення цього показника, що не дозволяє безпосередньо використовувати його в процесі моделювання. З урахуванням гіпотези про те, що розподіл доходів в суспільстві підпорядковується логарифмічно-нормальному закону [145], розроблено метод визначення доходів різних груп населення на підставі відомих даних про їх середній рівень і про ступінь нерівномірності розподілу доходів (коефіцієнт фондів) [154]. Коефіцієнт фондів визначається, як відношення доходів 10% найбагатших громадян до доходів 10% найбідніших. В Україні його величина в 2000-2015 роках коливалася в межах 5,5 – 7 одиниць.

Доходи служать основою формування *грошових накопичень*. Вважається, що частка витрат на накопичення становить до 15% доходів, однак цей показник в свою чергу залежить від багатьох інших факторів, серед яких, наприклад, рівень доходів і витрат домогосподарства, стабільність оточення, інфляційні очікування.

Максимальна сума кредиту, який банк може видати позичальнику, залежить від співвідношення кредитного платежу і доходу позичальника. Щомісячна плата за користування банківським кредитом в середньому не повинна перевищувати 50% місячного доходу особи, хоча до 2008 року деякі банки видавали кредит, навіть якщо сума платежу становила до 80% місячного доходу позичальника, що також слід враховувати.

Максимальну суму, яку може отримати клієнт, зручно розраховувати користуючись формулою визначення розміру ануїтетного платежу [7]:

$$A = \frac{S}{\left[\frac{(1+r)^n - 1}{(1+r)^n \cdot r} \right]}, \quad (3.16)$$

де S – сума кредиту; A – кредитний платіж; r – відсоткова ставка (за 1 період); n – кількість періодів.

Вираз у квадратних дужках називається «ануїтетний множник», і може використовуватися для визначення суми кредиту, якщо відомий допустимий розмір виплат.

$$S = \left[\frac{(1+r)^n - 1}{(1+r)^n \cdot r} \right] \cdot D \cdot cr, \quad (3.17)$$

де D – дохід позичальника; cr – припустима частина доходу, яка може бути спрямована на погашення кредиту.

Іпотечне кредитування українські банки почали пропонувати населенню приблизно в 2000 році. Спочатку терміни кредитів були короткими (до 3 років), а ставки – великими. До 2008 року термін кредитування поступово зростав, а ставки, навпаки, знижувалися. Хоча умови кредитування різних банків відрізнялися, в цілому ситуацію на ринку визначали декілька найбільших фінансових установ. Аналіз архівів банківських прес-релізів та інших подібних джерел дозволив в цілому встановити динаміку зміни умов кредитування на покупку нерухомості. Слід зазначити, що ціна на житлову нерухомість в Україні традиційно враховується в доларах США (USD). Причому якщо до 2008 року велика частина кредитів так і видавалися в доларах, з огляду на більш низькі процентні ставки, то після цього валютне кредитування фактично було заборонено і актуальними ставками по іпотеці стали гривневі.

На підставі моделей причинно-наслідкових взаємозв'язків, наведених на рис. 3.15 і 3.16, побудована динамічна імітаційна модель ціноутворення на ринку житлової нерухомості (рис. 3.17).

Основні параметри, які визначають фінансові можливості покупців нерухомості в м. Київ наведено в табл. 3.12 і табл. 3.13.

Таблиця 3.12 – Параметри моделі фінансових можливостей покупців в 2000-2007 роках

№ _{з/п}	Параметр:	Рік							
		2000	2001	2002	2003	2004	2005	2006	2007
1.	Дохід, 10 перцентиль, USD/рік	2012	2676	3306	3983	5049	7240	9619	12967
2.	Строк кредитування, років.	5	5	7	10	15	12	25	35
3.	Ставка, % річних, USD	35%	30%	16%	14%	13%	15%	13%	12%
4.	Ставка, % річних, грн.	-	-	25%	27%	12%	19%	17%	13%
5.	Потрібний рівень доходів	0,8	0,8	0,8	0,8	0,8	0,8	0,8	0,8

Таблиця 3.13 – Параметри моделі фінансових можливостей покупців в 2008-2015 роках

№ _{з/п}	Параметр:	Рік							
		2008	2009	2010	2011	2012	2013	2014	2015
1.	Дохід, 10 перцентиль, USD/рік	18088	11529	12516	14633	16992	18465	13070	9263
2.	Строк кредитування, років	25	10	8	10	10	15	10	10
3.	Ставка, % річних, USD	14%	-	-	-	-	-	-	-
4.	Ставка, % річних, грн.	18%	26%	25%	25%	24%	20%	25%	25%
5.	Потрібний рівень доходів	0,8	0,5	0,5	0,5	0,5	0,5	0,5	0,5

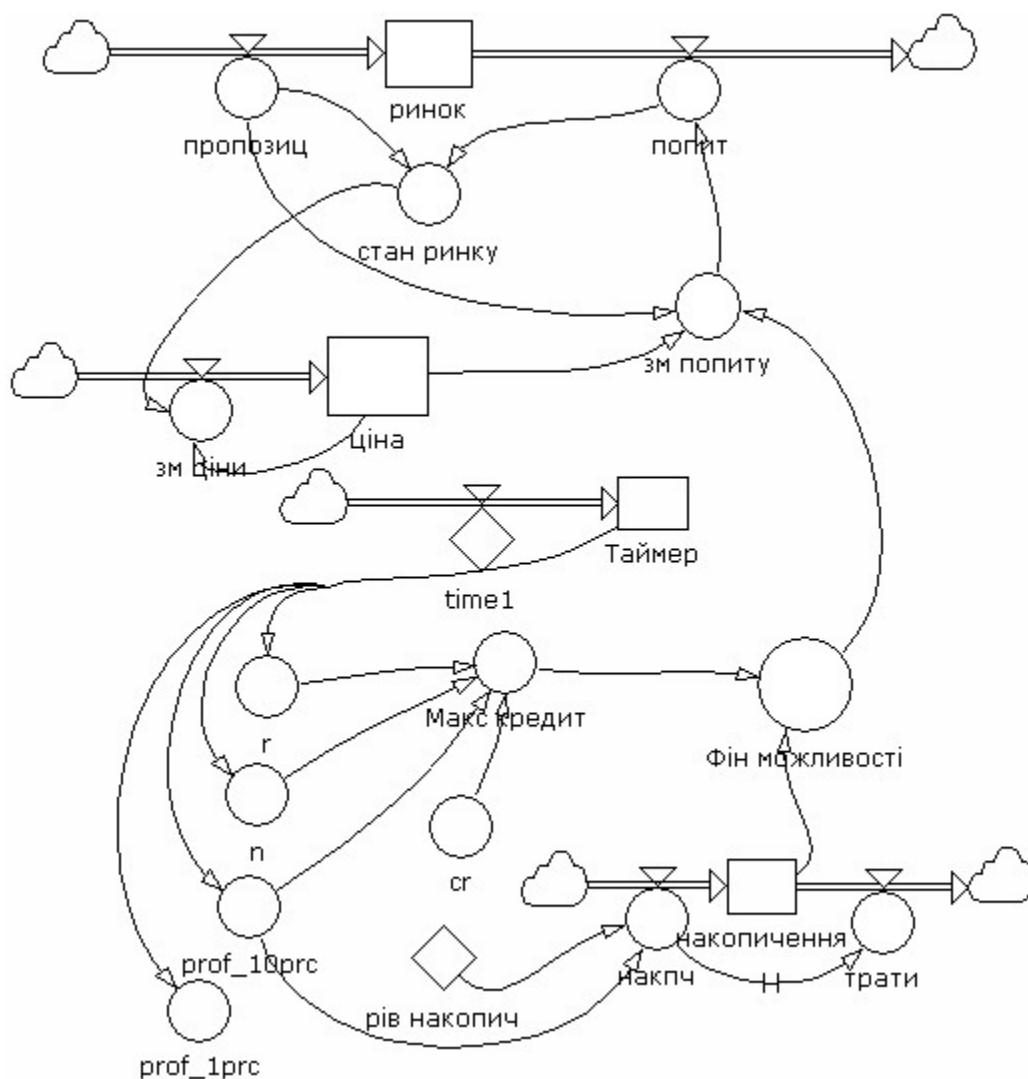


Рисунок 3.17 – Динамічна імітаційна модель ціноутворення на ринку нерухомості

Аналіз статистики продажів квартир в період ажіотажного попиту, показує, що в середньому пропозицію в м. Київ становить приблизно 3000 квартир на місяць. Однак, з огляду на вимоги гомогенності, в моделі прийнято значення 1000 в місяць, що відповідає обсягам продажу тільки однокімнатних квартир.

Змінна «стан ринку» в моделі задана з урахуванням несиметричної поведінки продавців, яке виражається в тому, що вони підвищують ціни охочіше, ніж знижують. Так, за одиницю системного часу максимальне збільшення ціни може скласти 10%, а зниження - 5%, що задано в системі *PowerSim* наступним чином:

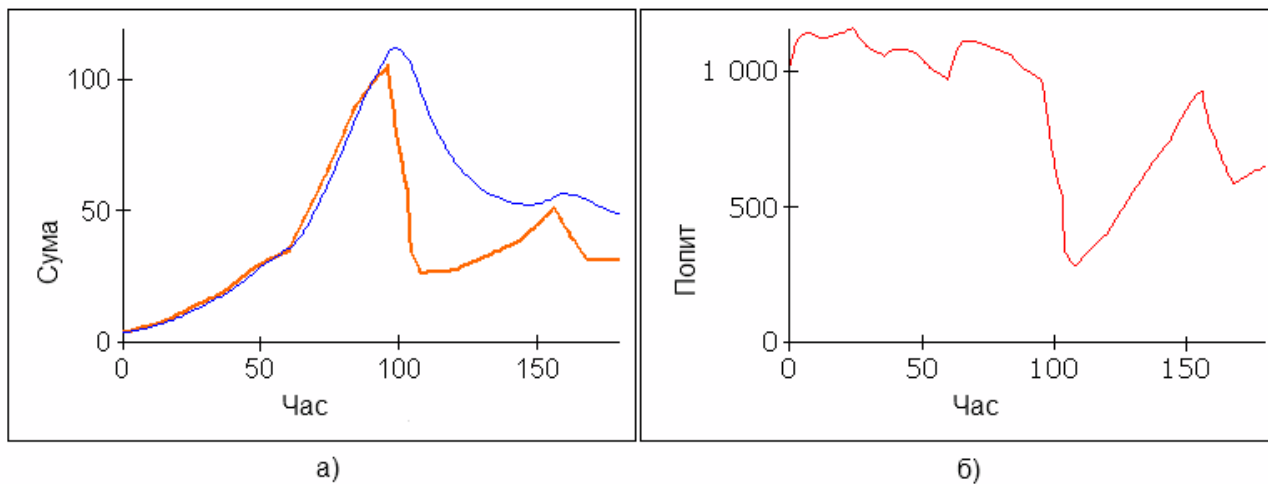
GRAPH(Попит/пропозиція;0;0,2;{0,95;0,96;0,98;0,99;1;1,02;1,06;1,09;1,1})

Розглянемо декілька імітаційних експериментів із даною моделлю.

Імітаційний експеримент I.3.1.

Мета: порівняти результати імітації, заснованої на фактичних даних про динаміку основних параметрів банківських кредитів з реальними змінами цін на нерухомість в Києві.

Результати імітаційного експерименту I.3.1 можна проілюструвати графіками, які показано на рис. 3.18. З графіків видно, що після 96 кроку (це відповідає 2008 року) різко падають фінансові можливості потенційних покупців житла (рис. 3.18а). Так само різко падає попит (рис. 3.18б). Однак ціна при цьому зменшується не так швидко і нижню межу проходить вже в той час, коли фінансові можливості покупців знову істотно зростають (рис 3.18а).



а) фінансові можливості покупця (потовщена лінія) і ціна на житло (тонка лінія), тис. USD; б) величина попиту на житло, шт.

Рисунок 3.18 – Динаміка зміни показників в імітаційному експерименті I.3.1

Але слід порівняти результати моделювання з реальними змінами цін на нерухомість в Києві. Для цього простежимо ціни на однокімнатні квартири в таких районах, як Дніпровський, Подільський і Святошинський, для яких фактор престижності не грає значної ролі. Результати зіставлення наведено на рис. 3.19.

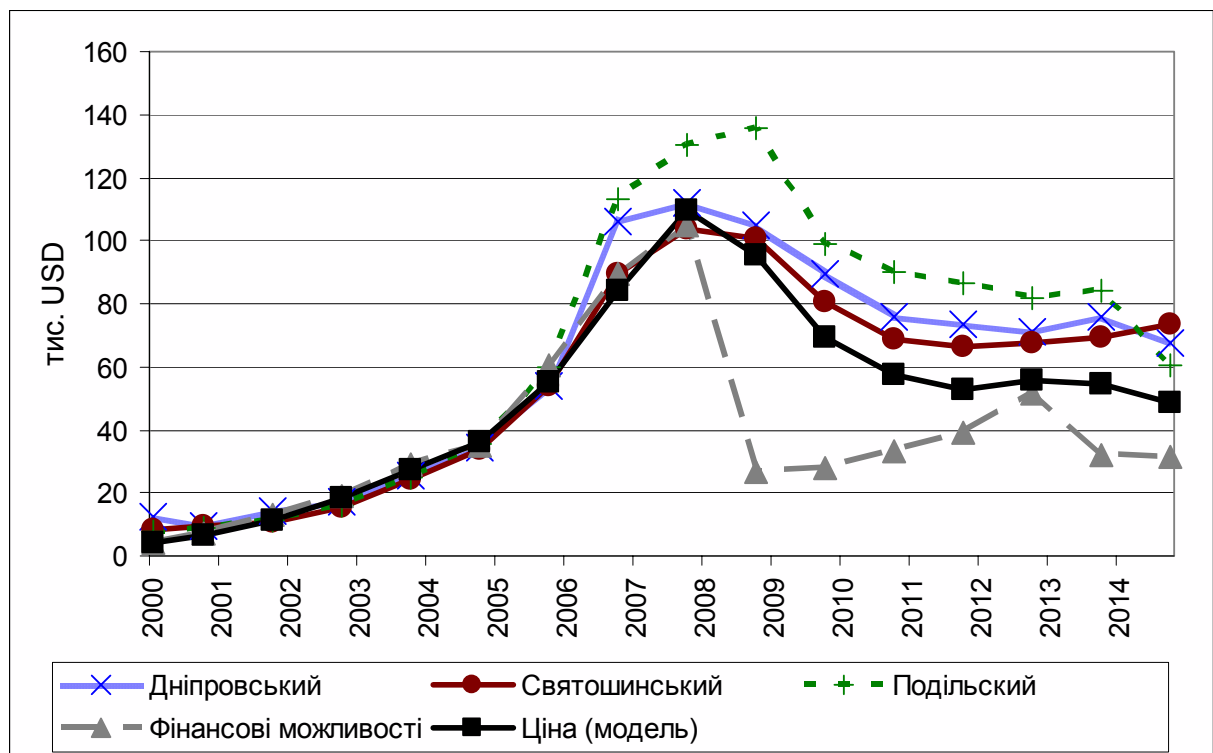


Рисунок 3.19 – Зіставлення результатів імітаційного експерименту з реальною динамікою цін на нерухомість в Києві в 2000-2015 рр. (тис. USD)

Як видно з графіка (рис. 3.19), результати імітаційного експерименту майже точно збігаються з фактичними даними на стадії зростання цін. Що ж стосується стадії падіння цін, то найбільш точний збіг вийшов із цінами на нерухомість в Дніпровському і Святошинському районах. Ціни ж в Подільському районі виявилися трохи вище прогнозованих значень. Очевидно, в цьому випадку на ціни впливають додаткові фактори, які не враховано в моделі (престижність, коливання пропозиції тощо).

Імітаційний експеримент І.3.2.

Мета: проаналізувати розвиток ринку житлової нерухомості, за умов, що банки взагалі не видають іпотечні кредити і на динаміку цін впливають тільки доходи населення.

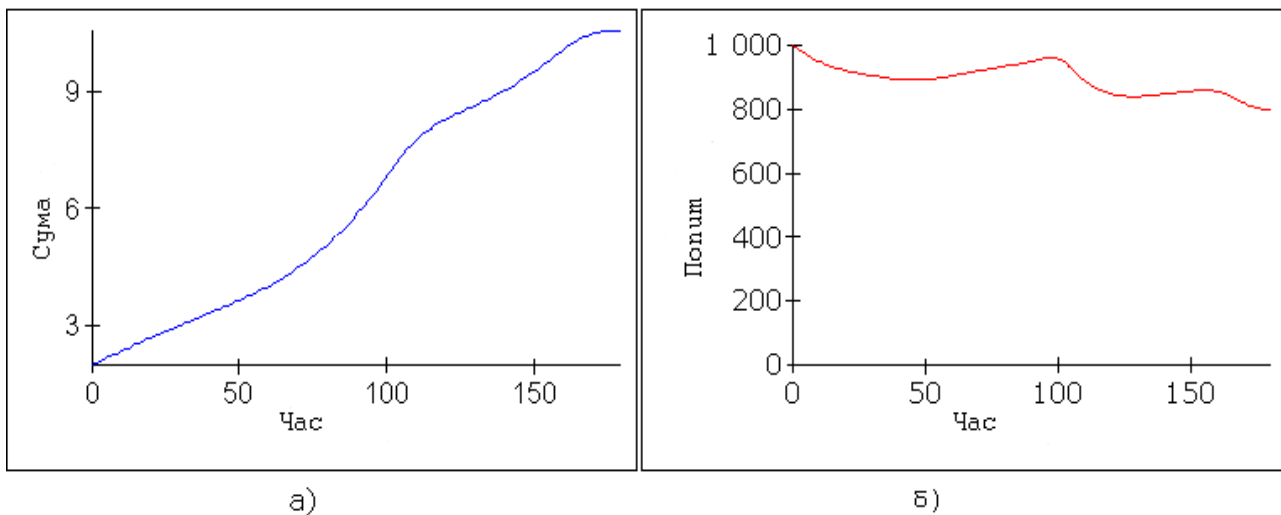
Результати експерименту І.3.2 показано на рис. 3.20 (час на графіку відповідає кількості місяців, що минуло з січня 2000 року). Як видно з рис. 3.20, в цьому випадку, ціни на однокімнатні квартири, незважаючи на кризи 2008 і 2014 років, росли б досить рівномірно і до теперішнього часу досягли б позначки близько 10 тис. USD.

Імітаційний експеримент І.3.3.

Мета: проаналізувати розвиток ринку житлової нерухомості, за умов, що кризові явища не торкнулися банківської сфери, і термін іпотечного кредитування в Україні досяг би 50 років, що є нормальним для розвинених країн Заходу.

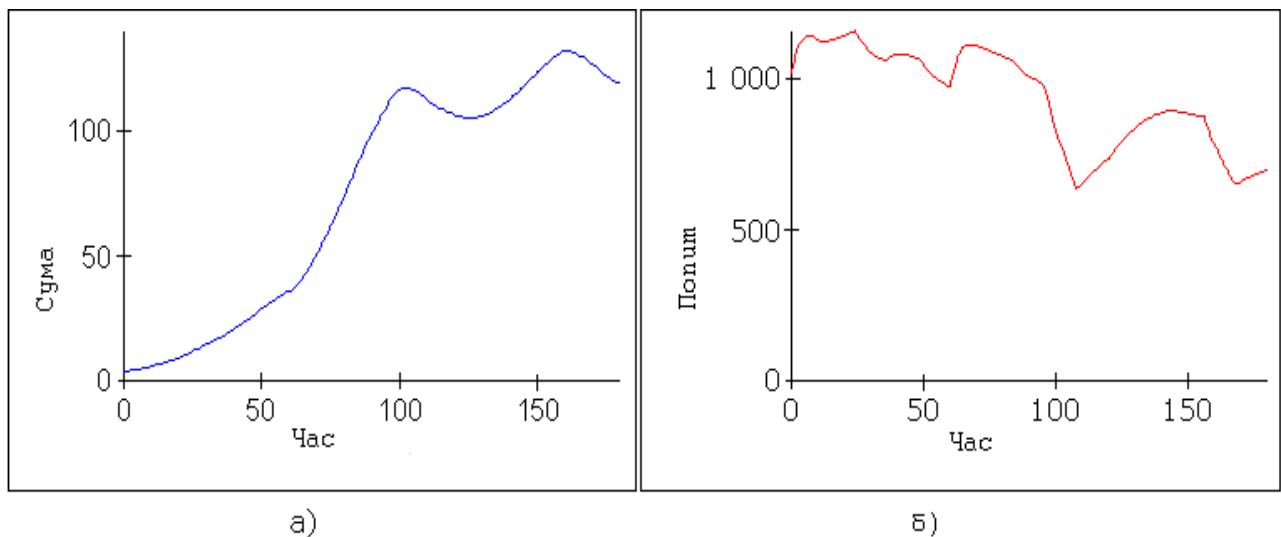
Результати експерименту І.3.3 показано на рис. 3.21. Як видно з рис. 3.21, і в цьому випадку нескінченного зростання не відбувається. Результати моделювання показують, що ціна на нерухомість стабілізується на рівні 120-

130 тис. USD за однокімнатну квартиру.



а) ціна на житло (тис. USD); б) величина попиту на житло (шт).

Рисунок 3.20 – Динаміка зміни показників в імітаційному експерименті І.3.2



а) ціна на житло (тис. USD); б) величина попиту на житло (шт).

Рисунок 3.21 – Динаміка зміни показників в імітаційному експерименті І.3.3

Слід звернути увагу на те, що динаміка доходів населення в імітаційних експериментах І.3.1-І.3.3 залишалася незмінною.

З проведених експериментів можна зробити висновок, що найсильнішим фактором ціноутворення на ринку нерухомості є зовнішні джерела фінансування, а саме – банківський кредит.

Результати моделювання мають як практичне, так і теоретичне значення.

Теоретичне значення результатів в тому, що запропонована модель дозволяє проаналізувати розвиток ринку в гіпотетичних ситуаціях, а також встановити рівень впливу зовнішніх факторів на ціни і чутливість моделі до їх змін.

Практичне значення полягає у можливості прогнозування реакції ринку на зовнішні події для вибору оптимальної стратегії участі в операціях на ньому.

Моделювання внутрішніх факторів ризику банківських активних операцій, на прикладі процесів реструктуризації кредитів. Одним з найтяжчих наслідків кризових явищ в економіці для українських банків слід вважати надмірне зростання простроченої кредитної заборгованості, обсяг якої в десятки разів перевищує норми, які прийняті в світовій практиці, та вимірюється вже сотнями мільярдів гривень. Попередження ризиків виникнення неплатежів за кредитними угодами, таким чином безумовно є актуальним завданням для банківської системи України.

Однією з головних причин виникнення неплатежів по кредитах слід вважати неадекватний облік банками можливостей позичальників – фізичних осіб. У зв'язку з цим запропоновано імітаційну модель фінансового стану позичальників – фізичних осіб в умовах кризи, яка дозволяє прогнозувати наслідки зміни балансу доходів і витрат позичальника, наприклад зниження заробітної плати, або збільшення виплат по кредиту [159]. Проте, в період стабілізації, перехід до якого очікується в середині 2017 року, згідно з заявою НБУ про перехід до завершальної стадії очищення банківської системи [192], більш актуальним завданням для банків стає розробка методологічного забезпечення процесів реструктуризації проблемної кредитної заборгованості, тобто зміни умов кредитних договорів з метою відновлення платежів клієнтів.

Розглянемо принципи, на яких ґрунтується моделювання процесів реструктуризації кредитів.

По перше, у разі великої суми позики (іпотека, або придбання автомобіля), обслуговування кредиту стосується не лише позичальника, а й його родини. Тому при розгляданні сімейного бюджету слід враховувати сукупні доходи та витрати всієї сім'ї.

По друге, витрати сімейного бюджету прямо пропорційно залежать від його доходів.

Прибуткову частину сімейного бюджету складає переважно заробітна плата членів родини. Склад витратної частини є більш різноманітним. Умовно витрати сім'ї можна поділити на обов'язкові (витрати першої необхідності, оплату житлово-комунальних послуг і обслуговування кредиту), витрати на накопичення та інші витрати, до яких у свою чергу можна віднести умовно-обов'язкові, тобто такі, які позичальник вважає неминучими для своєї сім'ї, та необов'язкові, до яких відносяться усі інші. Сума обов'язкових витрат не залежить від розміру прибуткової частини бюджету. Сума коштів, що залишилися після усіх витрат (обов'язкових та інших) відноситься на накопичення.

При збільшенні або зменшенні прибутків сімейного бюджету, витрати змінюються не одночасно, а поступово, що обумовлено наявністю інерційного механізму психологічної адаптації до нових умов. Очевидно, що параметри цього механізму індивідуальні, причому швидкість адаптації до збільшення прибутків і до їх зменшення в загальному випадку є різною.

Розглянемо імітаційну модель сімейного бюджету, засновану на використанні принципів системної динаміки Дж. Форрестера [231]. Модель дозволяє проаналізувати зміну грошових потоків сімейного бюджету при зміні

балансу прибутків і витрат (рис 3.22).

Основними потоками в даній моделі є доходи (*Profit*) і витрати (*Expenses*) бюджету. Параметри моделювання задаються за допомогою системи констант і початкових умов, основними з яких є наступні: *Starting_profit* – початкові доходи сім'ї; *Credit* – сума платежів по кредитах; *First_need* – витрати першої необхідності; *Other_need* – умовно-обов'язкові витрати; *Equity* – сума наявних накопичень; *Luxury* – необов'язкові витрати.

Робота механізму адаптації визначається константами *time_grow* і *time_fall*. Чим більше їх значення, тим повільніше відбувається пристосування витрат сімейного бюджету до збільшення, або зменшення доходів. У певний момент в моделі відбувається раптова зміна доходів (*Profit*).

Початкові значення змінних і параметрів моделі наведено в табл. 3.14. Цей набір можна розглядати, як характеристику умовного банківського позичальника до початку кризи.

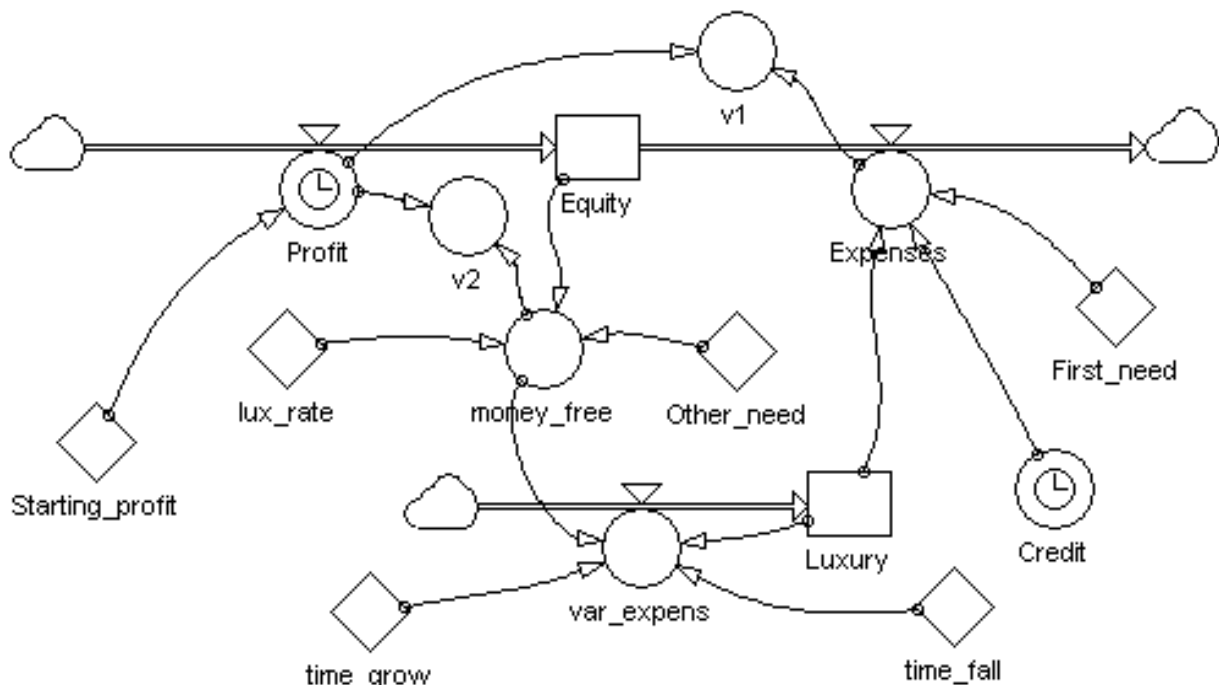


Рисунок 3.22 – Динамічна імітаційна модель сімейного бюджету, з урахуванням кредитних виплат

Таблиця 3.14 – Початкові значення змінних і параметрів моделі

№ _{пп}	Параметр		№ _{пп}	Параметр	
	Найменування	Значення		Найменування	Значення
1	Equity	2600	5	Starting_profit	3000
2	Luxury	1400	6	time_fall	6
3	Credit	1000	7	time_grow	4
4	First_need	600			

Розглянемо поведінку моделі сімейного бюджету при зменшенні доходів на 20%. Припустимо, таке зменшення станеться в 24-му періоді модельного

часу, що еквівалентно двом рокам в реальному масштабі. На рис. 3.23 показано динаміку змінних моделі в цьому випадку.

Аналіз рис. 3.23 дозволяє бачити, що при зниженні доходів, сума грошових коштів, вільних для фінансування інших витрат (змінна *Money_free*) зменшується, при тому, що самі витрати деякий час залишаються на колишньому рівні. Тому для їх покриття використовуються накопичені раніше грошові кошти (змінна *Equity*). Сума накопичень при цьому починає швидко знижуватися. Зменшення вільних грошових коштів спричиняє за собою поступове зменшення витрат, проте це зниження відбувається меншими темпами, ніж зменшення накопичень.

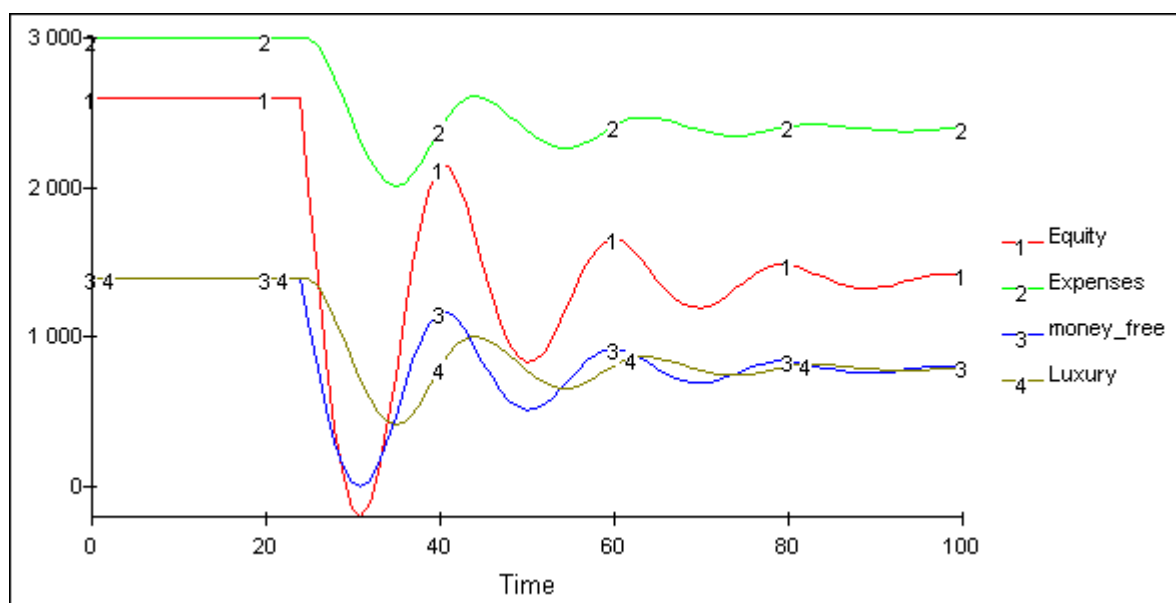


Рисунок 3.23 – Динаміка змінних моделі при $\Delta Profit = -20\%$

Як показує аналіз рис. 3.23, навіть невелика зміна доходів позичальника може привести до істотного падіння сальдо балансу сімейного бюджету, яке в даному випадку досягає значення -191 грн., тобто йде в область дефіциту. Але такий дефіцит у більшості випадків не призводить до катастрофічних наслідків, оскільки може бути порівняльне легко компенсований з накопичень, чи додаткових джерел доходу. Надалі коливання параметрів, що характеризують стан сімейного бюджету, поступово згасають і стабілізуються на нових стійких значеннях.

Очевидно, що в моделі динаміки сімейного бюджету, відновлення платоспроможності може настати завжди, якщо виконується співвідношення:

$$Profit \geq Credit + First_need + Other_need. \quad (3.18)$$

Однак, на практиці, при досить великому відносному падінні прибутків період неплатоспроможності збільшується настільки, що перевищує встановлені НБУ максимальні ліміти, внаслідок чого позичальник потрапляє в категорію безнадійних. Розглянемо, наприклад, результати моделювання при зниженні доходів не на 20%, а на 40%. (рис. 3.24 і табл. 3.15).

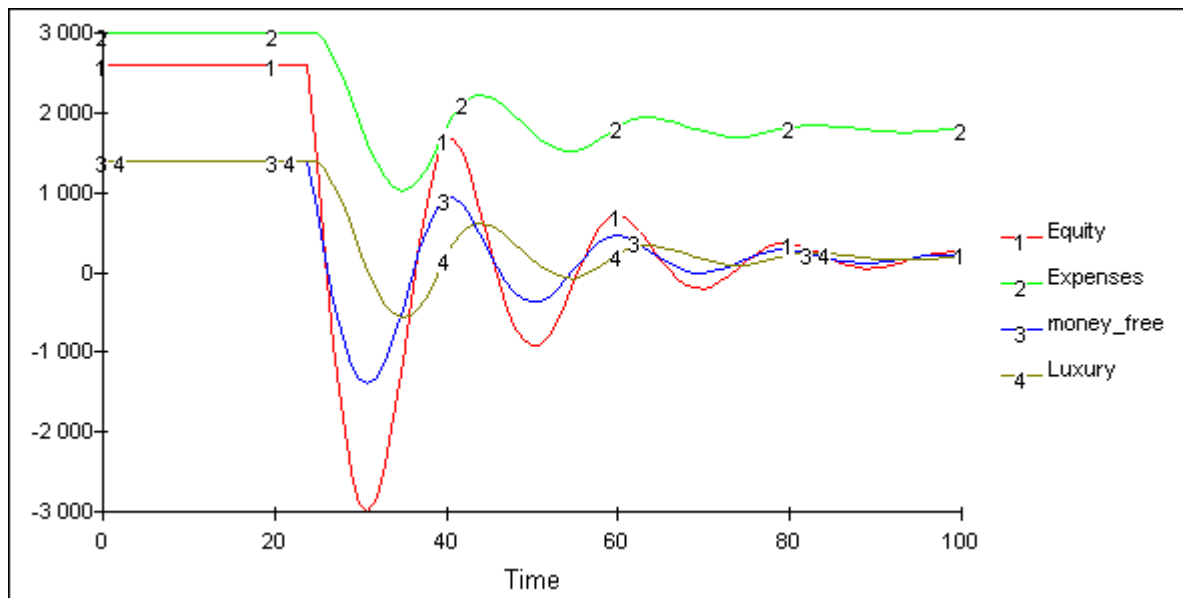


Рисунок 3.24 – Динаміка змінних моделі при $\Delta Profit = -40\%$

Таблиця 3.15 – Стан змінних моделі після впливу критичних збурень, при $\Delta Profit = -40\%$

Time	Profit	Credit	Equity	Luxury	Expenses
23	3 000,00	1 000,00	2 600,00	1 400,00	3 000,00
24	1 800,00	1 000,00	2 600,00	1 400,00	3 000,00
25	1 800,00	1 000,00	1 400,00	1 400,00	3 000,00
26	1 800,00	1 000,00	200,00	1 300,00	2 900,00
27	1 800,00	1 000,00	-900,00	1 116,67	2 716,67
28	1 800,00	1 000,00	-1 816,67	872,22	2 472,22
29	1 800,00	1 000,00	-2 488,89	592,13	2 192,13
30	1 800,00	1 000,00	-2 881,02	302,70	1 902,70
31	1 800,00	1 000,00	-2 983,72	28,83	1 628,83
32	1 800,00	1 000,00	-2 812,55	-207,95	1 392,05
33	1 800,00	1 000,00	-2 404,60	-391,00	1 209,00
34	1 800,00	1 000,00	-1 813,60	-509,55	1 090,45
35	1 800,00	1 000,00	-1 104,04	-559,09	1 040,91
36	1 800,00	1 000,00	-344,95	-532,33	1 067,67
37	1 800,00	1 000,00	387,38	-417,36	1 182,64
38	1 800,00	1 000,00	1 004,74	-239,60	1 360,40
39	1 800,00	1 000,00	1 444,34	-29,11	1 570,89
40	1 800,00	1 000,00	1 673,45	183,71	1 783,71

Аналіз табл. 3.15 та рис. 3.24 показує, що при зниженні прибуткової частини на 40% сімейний бюджет починає відчувати дефіцит вже через 3 місяці і тимчасово повертається до позитивного сальдо лише через 11 місяців після цього. В той же час, згідно з чинним законодавством, прострочення виплат по кредиту більш ніж 91 день служить підставою для віднесення його до категорії «безнадійних» [199].

Для зменшення боргового навантаження на позичальника комерційними банками використовуються різні методи реструктуризації умов кредитних договорів.

В практиці українських банків зустрічаються такі схеми реструктуризації [212, 176]:

- заміна валюти кредиту з іноземної на національну;
- заміна схеми нарахування відсотків з погашення основної частини боргу рівними частинами на ануїтет;
- збільшення терміну кредитного договору.

Недоліком перелічених схем є відносно невелике зменшення розміру платежу (як правило, не більше 20%), що забезпечується з їх використанням. Крім того можна виділити і інші недоліки, які характерні в цілому для банківської системи України, зокрема:

- запізнення початку процесу реструктуризації, внаслідок чого фінансовий стан позичальника погіршується настільки, що виконання умов нового договору для нього часто також є скрутним;
- реструктуризації піддаються тільки умови, які відносяться до виплати основної суми боргу, тоді як виплати відсотків по кредиту банки продовжують вимагати в повному обсязі;
- умови кредитного договору змінюються до кінця його дії, тобто навіть після того, як позичальник впорається зі своїми фінансовими утрудненнями, для нього продовжуватимуть діяти умови договору що реструктуризовано.

Перелічені недоліки не лише знижують ефективність реструктуризації, як методу зменшення фінансового навантаження, але і призводять до росту недовіри населення до банківської системи, зважаючи на наявність очевидного розриву між декларованою лояльністю до клієнта і фактичними діями банків.

Запропонована імітаційна модель фінансового стану позичальника (рис. 3.21) дозволяє зробити детальний аналіз наслідків застосування різних схем реструктуризації умов кредитування. Розглянемо детальніше випадок зниження доходів умовного позичальника на 40%.

Як впливає з табл. 3.15, вже на кроці 26 баланс прибутків і витрат позичальника знижується майже до нуля. На наступному кроці прибутки позичальника вже не можуть покривати його витрати. Очевидно, що цей період є критичним для ухвалення рішення про реструктуризацію заборгованості, оскільки інакше фінансовий стан клієнта може стати безнадійним. З іншого боку, ухвалення рішення про реструктуризацію умов кредиту раніше, ніж через місяць з моменту зниження доходів, важко здійснити, як за психологічної інерції клієнта, так і з організаційних причин. Отже, оптимальним часом для початку процесу реструктуризації слід прийняти термін 1-3 місяці, з моменту падіння доходів клієнта.

Розглянемо різні варіанти зниження платежу по кредиту і їх вплив на платоспроможність клієнта (табл. 3.16).

Таблиця 3.16 – Результати моделювання впливу різних умов зменшення платежу на стан сімейного бюджету

№пп	Зменшення платежу, %	Мінімальне сальдо бюджету, грн.	Довжина періоду із негативним сальдо, міс.
1	25%	-1960	9
2	50%	-1040	7
3	75%	-316	4

З аналізу даних, які наведено в табл. 3.16, випливає, що традиційні заходи по реструктуризації умов кредиту, які зазвичай дозволяють знизити платіжне навантаження на позичальника не більше, ніж на 25%, не дозволяють в достатній мірі нівелювати негативні наслідки падіння доходів позичальника. Тільки істотне зниження платіжного навантаження дозволяє йому вийти з кризи. Проте величина платежу не може постійно залишатися на такому рівні, оскільки не покриває навіть витрати банку на оплату грошових ресурсів, необхідних для кредитування. Тому, якщо умову (3.14) виконано, то після стабілізації фінансового стану позичальника необхідно поступово збільшити платіж до початкових значень. Якщо умову (3.14) не виконано, банк спільно із позичальником можуть переглянути умови кредитування у бік деякого зменшення щомісячного платежу із застосуванням однієї з існуючих схем реструктуризації [174].

Розглянемо різні варіанти поведінки моделі позичальника при збільшенні виплат по кредиту до початкового значення. В першому випадку платіж одночасно відновлюється до початкового значення. Кращим моментом для такого підвищення є мінімальне значення інших витрат сімейного бюджету (змінна *Luxury*) (рис. 3.25).

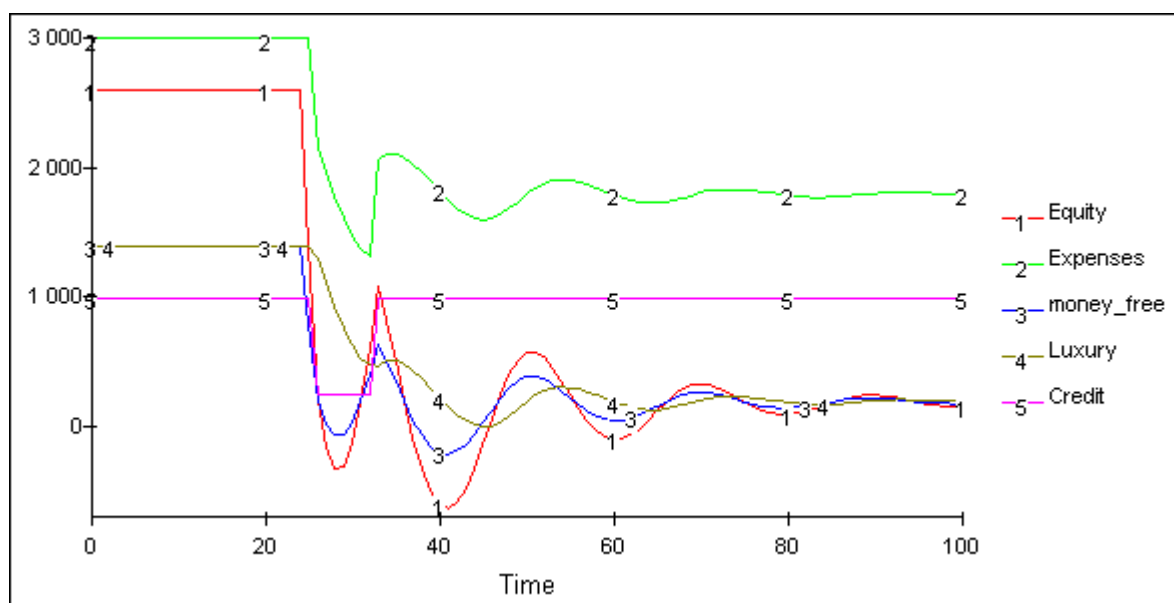


Рисунок 3.25 – Динаміка змінних моделі при одноступінчатому відновленні платежу

З аналізу рис. 3.25 можна побачити, що незважаючи на відновлення

розміру платежу до первинного рівня лише через 7 місяців, катастрофічного зниження балансу доходів/витрат позичальника не спостерігається. Так, мінімальне сальдо бюджету в цьому випадку складає -622 грн.

Головним недоліком такого методу є складність у визначенні моменту для відновлення розміру платежу, оскільки банк не може безпосередньо контролювати інші витрати сімейного бюджету. Кращим орієнтиром є сальдо сімейного бюджету (змінна *Equity*), яке банк може контролювати, зважаючи, наприклад, на своєчасність погашення позичальником кредитної заборгованості. Експерименти із моделлю довели, що в такому випадку слід використовувати двоступінчатую схему відновлення платежу (рис. 3.26).

В цієї схемі перше підвищення (до рівня 75% від первинного розміру) відбувається як тільки сальдо сімейного бюджету сягає позитивних значень, що в нашому прикладі відбувається через 5 місяців з початку проведення заходів щодо реструктуризації умов договору. Ще через 3-4 місяця платіж може бути відновлено до первинного розміру. У цьому випадку нове зниження сальдо сімейного бюджету виявляється незначним (-211 грн.), що навіть менш, ніж у попередньому прикладі із одноступінчатим відновленням розміру платежу.

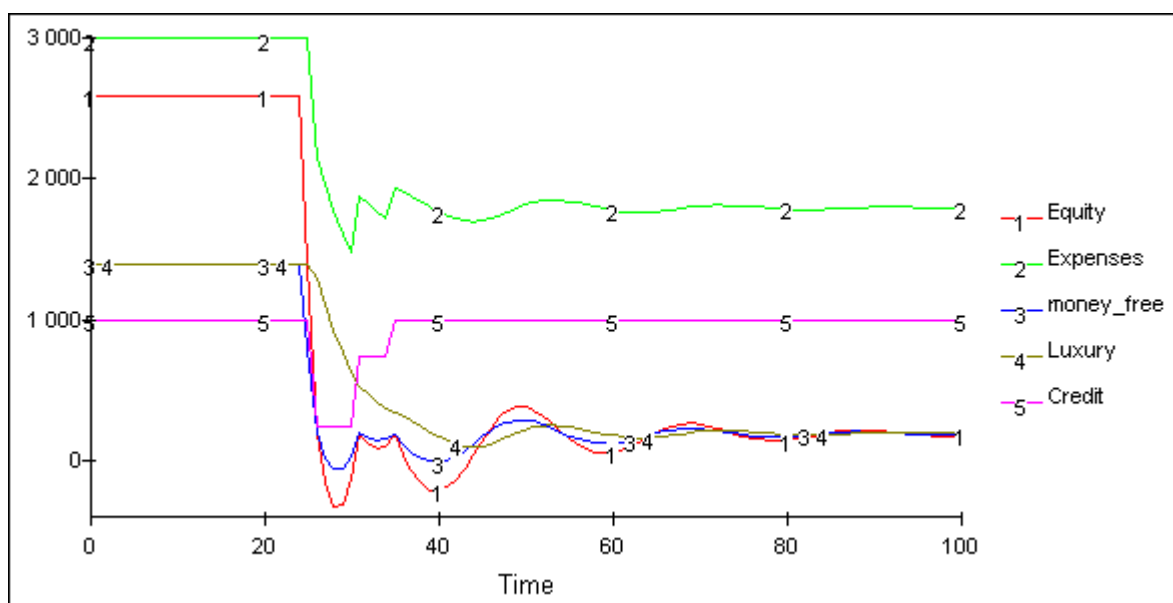


Рисунок 3.26 – Динаміка змінних моделі при двоступінчатому відновленні платежу

Очевидно, що для компенсації тимчасового зниження кредитного платежу загальний термін дії договору потрібно збільшити. Тому при практичному впровадженні запропонованого підходу однією з істотних проблем є узгодження з нормами банківського законодавства, що вимагають обов'язкового і своєчасного погашення відсотків по кредиту. Ця проблема може бути розв'язана як на державному рівні, через прийняття нормативних актів, що регламентують процеси реструктуризації кредитних договорів, так і на локальному рівні, через оформлення додаткових угод до кредитних договорів. У реальних умовах впровадження запропонованої схеми реструктуризації на основі динамічної імітаційної моделі сімейного бюджету дозволить уникнути

ризикі виникнення простроченої заборгованості за кредитними договорами.

Таким чином, за результатами проведених досліджень доведено твердження те, що ефективність вирішення складних економічних задач залежить від методів і інструментів інтелектуальних обчислень, які застосуються для цього, а також необхідність рішення економічних задач в різних постановках для пошуку найбільш ефективного інструменту моделювання. Так, для вирішення задачі біржового спекулянта більш ефективною виявилася регресійна постановка, а для задачі прогнозування банкрутств комерційних банків – кластеризаційна.

Досліджено використання генетичних алгоритмів для вирішення задачі спрощення динамічних рядів, яка відноситься до задач обробки даних. Запропонований метод дозволяє ефективно скорочувати кількість точок відліку ряду з будь-яким ступенем стиснення даних, та має такі переваги перед існуючими методами, як то збереження пікових значень ряду, а також більш ефективне стиснення даних, якщо значення в ряду змінюються повільно.

На прикладах аналізу зовнішніх та внутрішніх факторів кредитного ризику комерційних банків доведено ефективність використання методів системно-динамічного імітаційного моделювання для непрямой оцінки параметрів економічних систем в задачах прийняття рішень. Розроблено динамічну імітаційну модель ціноутворення на ринку житлової нерухомості, яка дозволяє оцінювати зміни цін на вторинному ринку житла, в залежності від зовнішніх факторів. Розроблено динамічну імітаційну модель сімейного бюджету кредитопозичальника та методи її застосування при аналізі процесів реструктуризації кредитів.

РОЗДІЛ 4. РЕАЛІЗАЦІЯ ІННОВАЦІЙНИХ МОДЕЛЕЙ І МЕТОДІВ ПРИЙНЯТТЯ РІШЕНЬ В ЕКОНОМІЧНИХ ЗАДАЧАХ

4.1 Вибір інструментальних засобів розробки інтелектуальних систем прийняття рішень

Аналіз програмного забезпечення для моделювання штучних нейронних мереж.

Незважаючи на існування «справжніх» нейрокомп'ютерів, в яких нейрони, синаптичні зв'язки між ними, а часто і алгоритми, навчання реалізовані апаратно [110], в економіці частіше знаходять застосування більш універсальні і дешеві рішення, засновані на програмній емуляції роботи штучних нейронних мереж. Це дозволяє отримати всі переваги нейрокомп'ютерів на платформі звичайних персональних комп'ютерів. У порівнянні з апаратною реалізацією, такий підхід має незрівнянно більшу гнучкість щодо використання різних архітектур і алгоритмів навчання. Його поширенню сприяє також розвиток технологій кластерних і хмарних обчислень, а також можливість використання для математичних обчислень процесорів, розташованих в сучасних графічних прискорювачах [92, 108].

Розглянемо деякі з програмних продуктів для моделювання ШНМ, представлених в даний час на ринку.

Серед *програмного забезпечення низькорівневої розробки* найбільш популярними мовами для програмування задач, пов'язаних з інтелектуальними обчисленнями є Python, R, C++. Розглянемо їх.

Python – поширена мова програмування, що відрізняється доброю читаністю коду і високою швидкістю процесу розробки програмного забезпечення. Основні реалізації мови поширюються вільно. У сукупності з архітектурними особливостями мови, які добре підходять для математичних обчислень, це виділяє Python, як одну з основних мов програмування наукових розрахунків. Для Python створена велика кількість бібліотек, що містять функції аналізу і обробки даних, що наближає його до програмованих систем математичного моделювання. Оскільки Python – мова програмування, яка інтерпретується, її недоліком є порівняно низька швидкість роботи. Однак вона використовується для створення широкого класу програмних продуктів, що спрощує інтеграцію інтелектуальних обчислень до існуючих систем [128].

R – мова програмування, яка цілеспрямовано орієнтована на виконання статистичних розрахунків. Тому вона є ефективною для вирішення задач аналізу і обробки даних. Для R існують як безкоштовні реалізації, так і комерційні. Останні містять вбудовані засоби для роботи з великими обсягами даних і інші розширення, які відсутні в базовій версії. Недоліками мови є її порівняно вузька спеціалізація, що ускладнює інтеграцію з іншими системами,

а також менший, у порівнянні з конкурентами, вибір бібліотек функцій [141].

C++ є однією з найпоширеніших мов програмування в світі. Для неї можна відзначити високу швидкість, великі обсяги напрацьованого програмного забезпечення, а також поширеність серед програмістів. Більшість реалізацій мови є комерційними, хоча є і вільно-поширювані версії. Серед недоліків також можна відзначити складний синтаксис і специфікацію мови, що збільшує поріг входження і перешкоджає його використанню нефахівцями [210]. Але відзначимо, що на даний час для C++ створено велику кількість вільно-поширюваних програмних каркасів і бібліотек, які реалізують основні алгоритми інтелектуальних обчислень, що істотно прискорює процес розробки [9].

Головною перевагою використання програмного забезпечення низькорівневої розробки є можливість реалізації будь-якої нестандартної архітектури ШНМ.

В клас *програмованих систем математичного моделювання* входять такі пакети програм, як Matlab, Octave, Scilab, Julia. Незважаючи на те, що Matlab - єдина комерційна розробка з перелічених програм, ця система є найбільш поширеною і фактично може розглядати як стандарт для даного класу. Структура і синтаксис мови програмування Matlab досить прості для освоєння, що забезпечує низький поріг входження. До складу системи входить велика кількість пакетів розширення, які роблять його універсальним засобом інженерних і наукових розрахунків. Документація до Matlab детальна, добре структурована і містить велику кількість прикладів. Також слід зазначити широкі можливості графічного представлення результатів.

До недавнього часу офіційне використання Matlab обходилося дорого, оскільки крім основного пакета необхідно було доплачувати окремо за кожне розширення. Однак зараз цінова політика компанії Mathworks змінилася в бік більшої доступності Matlab для приватного використання і освітніх цілей. Так, студентська ліцензія на пакет програм для аналізу даних коштує всього \$ 55, що майже в 300 разів менше їх вартості в складі комерційного продукту [183].

Решта продуктів класу *програмованих систем математичного моделювання* створювалися в значній мірі, як вільно-розповсюджувана альтернатива Matlab і мають схожий синтаксис мови і ідеологію. Самим близьким до Matlab і розвиненим пакетом є Octave, однак ця система гірше документована та має меншу кількість розширень [184].

Інтерактивні програмні платформи нейромережевого моделювання представлені найбільшою кількістю продуктів [10, 11]. Детальний їх розгляд виходить за рамки даного дослідження, тому обмежимося декількома типовими представниками цього класу програм.

Statistica Neural Networks компанії StatSoft. У складі Statistica присутні більшість відомих методів аналізу даних. Недоліком пакету є відносна складність освоєння і висока вартість. Проте, існує система знижок для освітніх установ [185].

Пакет *Neuro Solutions* відрізняється специфічним інтерфейсом і засобами візуалізації. Разом з тим пакет надає можливості з проектування нейронних

мереж нестандартних конфігурацій, різноманітні засоби обробки вхідних даних основних типів, включаючи графічну інформацію, засоби створення зовнішніх модулів, що розширюють базові можливості пакету [186].

Так як одним з перспективних напрямків застосування штучних нейронних мереж є використання їх для створення систем автоматичної торгівлі на валютних і фондових ринках. Це призвело до створення спеціалізованого пакета для аналізу фінансових часових рядів *NeuroShell Trader*. Можливості пакета дозволяють також використовувати його і для вирішення загальноекономічних задач, проте за зручністю використання і можливостям *NeuroShell* в цьому випадку поступається більшості програмних продуктів інших виробників з тієї ж цінової категорії [187].

Neural Designer – сучасний універсальний нейромережевий емулятор. Містить розвинені засоби аналізу і обробки даних, підбору архітектури ШНМ і аналізу результатів, що сприяє його застосуванню для вирішення економічних задач. Один з небагатьох продуктів, що підтримують можливість глибинного навчання нейронних мереж. *Neural Designer* є комерційним продуктом, проте з березня 2017 року доступна повнофункціональна безкоштовна версія без обмеження по терміну використання, але лімітована за обсягом вхідних даних (5000 рядків) [188].

Оскільки нейронні мережі широко застосовуються для пошуку і аналізу прихованих залежностей в даних, функції створення та аналізу ШНМ включено до складу більшості комерційних продуктів цього напрямку. Серед них можна виділити пакет *Deductor Studio*, розроблений компанією *BaseGroup*. Вартість аналітичної платформи *Deductor* нижче аналогів, при еквівалентних можливостях. До особливих рис продукту слід віднести розвинені засоби аналізу адекватності моделей [189].

Більшість безкоштовних інтерактивних програмних платформ нейромережевого моделювання слабо придатні для практичного використання при розробці інтелектуальних систем прийняття рішень в економіці. В основному такі програми створюються або для вирішення вузькоспеціалізованих завдань, або для використання в навчальних цілях [11].

Оскільки вартість є одним з основних критеріїв вибору інструментальних засобів розробки інтелектуальних систем прийняття рішень, порівняємо за цим критерієм продукти, що відносяться до категорій програмованих систем математичного моделювання і інтерактивних програмних платформ (табл. 4.1).

З аналізу табл. 4.1 можна зробити такий висновок:

Найнижча вартість комерційної версії у *Deductor Studio*. Серед можливостей, які надає ця система можна виділити:

1. Операції очищення даних (відновлення і згладжування, факторний аналіз, кореляційний аналіз, усунення дублікатів).

2. Операції, пов'язані з трансформацією даних (таблична обробка, перетворення ковзаючим вікном, сортування, застосування користувацьких формул і інші).

3. Пошук залежностей в даних – Data Mining (автокореляція, лінійна і логістична регресія, нейронні мережі, дерева рішень, карти Кохонена,

асоціативні зв'язки, кластеризація).

Таким чином, функції, які реалізовано в Deductor Studio, охоплюють більшість етапів циклу розробки ІСПР, включаючи попередню обробку даних, побудову моделі і аналіз результатів. Однак штучні нейронні мережі в рамках цього пакету є лише одним з методів аналізу даних, що обумовлює реалізацію тільки двох їх видів (персептрони і мережі Кохонена) та основних алгоритмів навчання.

Таблиця 4.1 – Порівняння вартості комерційних програмних пакетів нейромережевого моделювання (на 2017 р.)

№ з.п.	Назва продукту	Версії продукту		Обмеження для освітньої версії
		Комерційна	Освітня	
1	Matlab	\$2350 + \$1200	от \$55	Відсутні засоби інтеграції
2	Statistica	от \$3400	\$160	—
3	Neuro Solutions	\$995 - \$3995	\$295	—
4	NeuroShell Trader	\$1495 - \$3495	—	—
5	Deductor Studio	\$680	\$0	Обмеження на експорт та імпорт даних
6	Neural Designer	\$2500	\$0	5000 строк в навчальній вибірці

Серед програмних продуктів, що мають більш широкі можливості в частині моделювання ШНМ, низькою вартістю і наявністю безкоштовної освітньої версії виділяється пакет Neural Designer. На відміну від інших інтерактивних програмних платформ, він також дозволяє використовувати глибинне навчання багатошарових нейронних мереж, що є актуальним при вирішенні задач, які пов'язані з пошуком складних взаємозв'язків в даних.

Але, незважаючи на важливість критерію вартості програмного забезпечення, він рідко є єдиним і, для розробки ІСПР вибір необхідно робити за сукупністю критеріїв різної значимості.

Розглянемо процедуру нечітко-логічного зіставлення і вибору інструментальних засобів розробки ІСПР. Для обмеження кола задач доведемо наступне твердження.

Твердження 4.1.

Нечітко-логічне зіставлення і вибір інструментальних засобів вирішення економічних задач доцільно проводити для таких фаз циклу розробки і реалізації ІСПР: попереднього аналізу даних, підготовки даних, моделювання і оцінки результатів.

Доведення твердження 4.1.

Процедури аналізу предметної області, попереднього збору даних і впровадження не пов'язані безпосередньо з інтелектуальними обчисленнями і

можуть бути реалізовані на базі інформаційних систем, наявних на об'єкті дослідження.

Фази попереднього аналізу даних і моделювання припускають роботу з великою кількістю алгоритмів, методів і моделей, вибір інструментів реалізації яких обумовлений особливостями об'єкта дослідження. З погляду інформаційних потоків з ними тісно пов'язані фази підготовки даних і оцінки результатів. Тому для реалізації вказаних фаз слід використовувати один і той же програмний продукт, а велика кількість зовнішніх факторів потенційних варіантів вибору обумовлюють доцільність використання нечітко-логічних методів вибору.

На етапі реалізації до програмного забезпечення додатково висуваються вимоги щодо інтегрованості, що обмежує вибір і не дає його зробити *a priori*, без урахування особливостей існуючої інформаційної системи.

Твердження 4.1 доведено.

Для порівняння візьмемо програмні продукти, які є представниками всіх трьох підходів до реалізації інтелектуальних обчислень:

Програмне забезпечення низькорівневої розробки: мови програмування C++, R, Python з необхідними бібліотеками.

Програмовані системи математичного моделювання: система Matlab с нейромережевим пакетом розширення.

Інтерактивні програмні платформи: аналітична платформа Deductor і нейромережевий емулятор Neural Designer.

У табл. 4.2. представлено результати експертного оцінювання інструментальних засобів для попереднього аналізу даних і моделювання за п'ятибальною шкалою відповідно до критеріїв K.2.2.1-K2.2.7 на умовному прикладі.

Таблиця 4.2 – Бальні оцінки досліджуваних програмних продуктів

№ пп	Критерии	C++	MatLab	Deductor Studio	Neural Designer	R	Python
1	Простота використання	1	3	5	4	2	2
2	Функціональна придатність	1	5	3	4	4	2
3	Часова ефективність	5	3	2	3	5	4
4	Вартість	3	2	3	2	4	5
5	Ресурсомісткість	5	4	4	3	4	4
6	Доступність використання	5	3	5	3	3	5
7	Мобільність	5	5	5	3	4	5
8	Загалом	25	25	27	22	26	27

Відповідність між бальними оцінками і нечіткими параметрами приналежності задається відповідно до табл. 4.3.

Таблиця 4.3 – Параметри приналежності оцінок привабливості інструментальних засобів розробки ІСПР у вигляді балів і лінгвістичних змінних

Інтервал значень	Кількість балів	Лінгвістична оцінка привабливості продукту
[0; 0; 0.2; 0.25]	1	Дуже низька
[0.15; 0.25; 0.35; 0.4]	2	Низька
[0.3; 0.4; 0.55; 0.6]	3	Середня
[0.45; 0.6; 0.7; 0.75]	4	Висока
[0.65; 0.75; 1; 1]	5	Дуже висока

Значущість різних критеріїв оцінки для фаз попереднього аналізу даних і моделювання задається в лінгвістичних змінних відповідно до табл. 4.4.

Таблиця 4.4 – Параметри приналежності значущості критеріїв оцінки інструментальних засобів розробки ІСПР

Інтервал значень	Лінгвістична оцінка значимості критерію	Критерії
[0; 0; 0; 0.4]	Не має значення	5. Ресурсомісткість 7. Мобільність
[0.2; 0.35; 0.5; 0.6]	Середня	3. Часова ефективність 6. Доступність використання
[0.45; 0.6; 0.7; 0.8]	Висока	4. Вартість 1. Простота використання
[0.65; 0.75; 1; 1]	Дуже висока	2. Функціональна придатність

Програмна реалізація нечіткої моделі оцінки інструментальних засобів розробки ІСПР проведена в середовищі Matlab. Текст програми наведено в додатку А.

Результати аналізу формуються у вигляді функцій приналежності оцінок, графічне представлення яких показано на рис. 4.1.

Аналіз графіків, які наведено на рис. 4.1, дозволяє зробити висновки про те, що в заданому полі вхідних даних та критеріїв найкращим програмним продуктом є Deductor Studio. Виконання попереднього аналізу даних і моделювання з використанням C і Neural Designer є недоцільним. Оцінки інших продуктів досить близькі, тому графічний аналіз не дає змогу зробити обґрунтовані висновки про їх взаємне розташування. В такому випадку слід скористатися результатами дефазифікації (табл. 4.5).

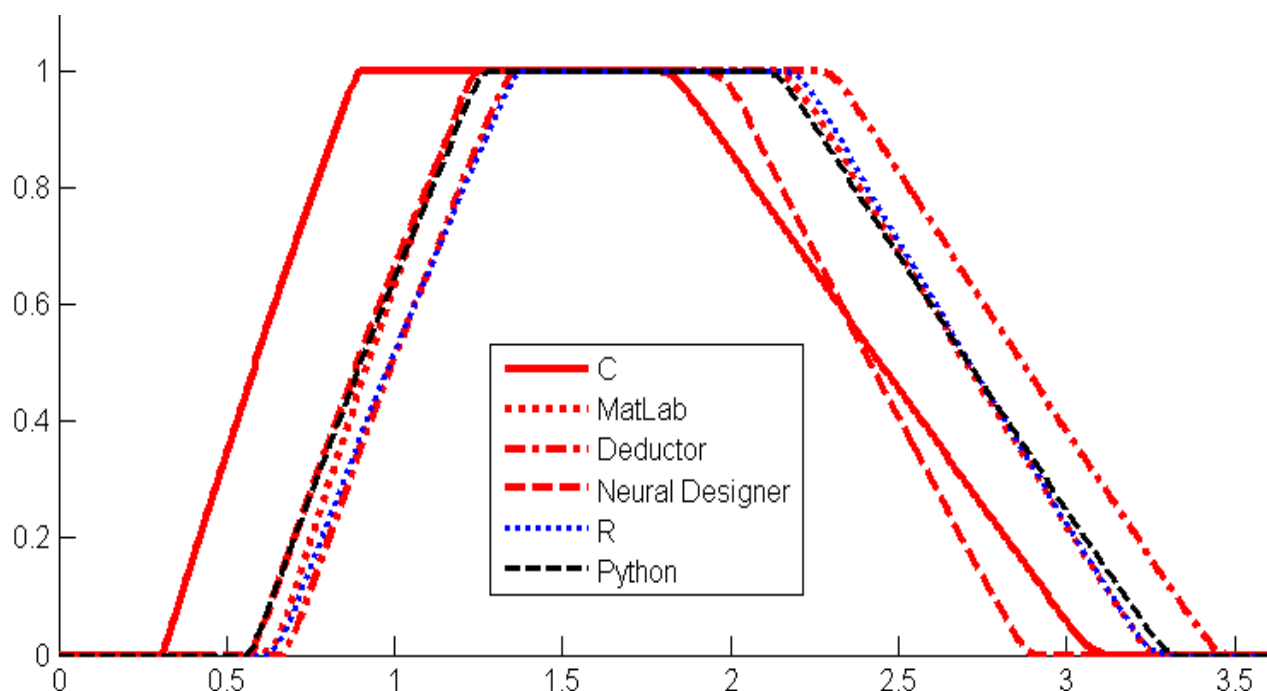


Рисунок 4.1 – Функції належності оцінок інструментальних засобів розробки ІСПР

Таблиця 4.5 – Дефазифіковані результати оцінки досліджуваних програмних продуктів

№ _{пп}	Програмний продукт	Дефазифікована оцінка	Місце
1	C++	1.5468	6
2	MatLab	1.8343	3
3	Deductor Studio	1.9545	1
4	Neural Designer	1.6693	5
5	R	1.8693	2
6	Python	1.8291	4
7	Абсолютно гірший	0.4610	–
8	Абсолютно кращий	2.8117	–

Для поліпшення порівнянності результатів в табл. 4.5 показані також оцінки умовних програмних продуктів, які оцінено за всіма критеріями в 1 бал (абсолютно найгірший) та 5 балів (абсолютно кращий).

Як видно з табл. 4.5, друге місце за привабливістю для вирішення завдань попереднього аналізу даних і моделювання займає мова програмування R, третє – програмована система математичного моделювання Matlab, четверте – мова програмування Python. Однак їх оцінки настільки близькі, що навіть при невеликих змінах значущості критеріїв, або бальних оцінок результати можуть змінитися. Тому на практиці слід вважати їх приблизно рівноцінними.

Аналіз програмного забезпечення для генетичного моделювання. Штучні нейронні мережі можуть розглядатися дослідником у вигляді свого роду «чорної скрині», внутрішня структура якої будується автоматично зі

стандартних елементів. Цим пояснюється наявність великої кількості програмних продуктів для їх моделювання. При розробці генетичних алгоритмів такий підхід неможливий з огляду на те, що найважливіший параметр ГА – функція пристосованості – є суто індивідуальною та повинна бути окремо запрограмована для кожної задачі. Внаслідок цього повністю інтерактивних пакетів програм для моделювання генетичних алгоритмів не існує, а більшість з існуючих продуктів являють собою набір функцій, реалізованих на якій-небудь мові програмування.

Аналіз програмного забезпечення з моделювання генетичних алгоритмів дозволяє класифікувати його на три основні категорії:

1. Низькорівневе;
2. Спеціальне;
3. Вбудоване.

Низькорівневе програмне забезпечення поширюється у вигляді набору процедур, що реалізують функції генетичних алгоритмів в рамках мов програмування широкого призначення. Перевагою такого програмного забезпечення є безкоштовність. Основним недоліком його використання є необхідність детального опису генетичної моделі аж до операцій кодування і декодування хромосом. Крім цього розробнику потрібно володіти навичками програмування на тій мові, на якій написаний пакет процедур. Так, для використання пакету генетичних процедур *Sugal 2.1* потрібно встановити компілятор Borland C++ версії 3.0, [137]. Ще одним недоліком можна вважати те, що в процесі написання програми потрібно не тільки описати саму модель, але і запрограмувати допоміжні функції введення і виведення інформації, що збільшує час розробки.

До *спеціального* програмного забезпечення можна віднести реалізації генетичних алгоритмів в рамках запрограмованих систем математичного моделювання, з яких найбільш поширеною є MatLab, до складу якої інструментарій розробки генетичних алгоритмів входить, починаючи з версії MatLab 7.0. Процес написання програм в системі MatLab значно простіше, ніж при використанні мов загального призначення, що пояснюється простою структурою мови, її орієнтацією на матричне представлення інформації, розвиненими можливостями з обробки даних та їх візуалізації. Пакет генетичного моделювання *Genetic Algorithm and Direct Search Toolbox*, що входить в MatLab, дозволяє при моделюванні здійснювати найбільш складні операції з кодування/декодування хромосом автоматично, що різко знижує трудомісткість розробки.

До категорії *вбудованого* віднесемо ПО, що створюється у вигляді доповнень до інших програмних продуктів, наприклад, до табличного процесору Microsoft Excel. Режим роботи при цьому є наближеним до інтерактивного. Від користувача потрібно ввести початкові дані, цільову функцію та обмеження моделі. Також можна відкоригувати деякі параметри генетичного алгоритму, після чого запустити оптимізацію. Оскільки принципи і навички роботи з Microsoft Excel відомі широкому колу користувачів, програмне забезпечення, що вбудовується, є найпростішим в освоєнні. На

ринку представлено кілька програмних продуктів цієї категорії, зокрема пакет *GeneHunter for Excel* компанії *WardSystems* і пакет *Evolver* компанії *Palisade Corporation*. Обидва продукти побудовані з використанням однієї ідеології і надають користувачам приблизно однаковий сервіс.

Щоб визначити область застосування кожної з зазначених категорій програмного забезпечення проведемо порівняльний аналіз відповідних продуктів за кількома критеріями: інтерфейс і зручність використання, швидкість роботи і вартість. Для розгляду візьмемо такі продукти – Sugal v.2.1, MatLab v.7.7 і Evolver v.6.0.

Пакет *Sugal* був розроблений в Великобританії в Університеті м. Сандерленда. Він реалізує всі основні функції генетичних алгоритмів. *Sugal* розрахований на роботу з компілятором Borland C++ версії 3.1., але може бути адаптований і до сучасних версій цієї мови програмування. Використання C++ дає можливість створювати готові до виконання програми, як для пакетного, так і для діалогового режиму. З переваг *Sugal* необхідно відзначити великі можливості по налаштуванню генетичних алгоритмів (близько 90 параметрів), зручний інтерфейс при роботі з уже готовою програмою, можливість виводити результати роботи в діалоговому режимі [138].

Інтерфейс *Sugal* дає доступ до налаштування параметрів алгоритму та інших функцій програми, а також дозволяє відображати графік роботи алгоритму. Наявність інтерфейсу діалогового режиму вигідно виділяє *Sugal* серед інших реалізацій генетичних алгоритмів на низькому рівні, хоча програмування ГА залишається досить трудомістким. Хромосома в *Sugal* розглядається як одновимірний масив елементів, на підставі якого функція користувача повинна розрахувати рівень пристосованості. При реалізації моделі необхідно визначити структуру хромосоми і скласти на мові C++ програму з визначення її пристосованості. Це значно збільшує трудомісткість, але дозволяє отримати найбільшу гнучкість моделей, а також оптимізує швидкість виконання програми.

Одним з недоліків *Sugal* є те, що в цьому пакеті не реалізовано сучасні модифікації генетичних алгоритмів, які покращують збіжність, дозволяють ефективно досліджувати безперервні функції і ряд інших. Однак програмний код пакета є відкритим, що залишає принципову можливість додавання нових функцій.

Система *MatLab* розробляється і випускається з 1984 року. Вона складається з базового пакету та розширень, які реалізують додаткові функції. Можливість моделювання ГА присутня в *MatLab* починаючи з версії 7.0. (2004 рік). Оскільки компанія MathWorks регулярно випускає нові версії системи, в ній присутня реалізація всіх новітніх розробок в області генетичних алгоритмів. Робота з *MatLab* можлива як в командному, так і в програмному режимах. Мова програмування набагато простіше в освоєнні, ніж мова C++, яка використовується в *Sugal*. Оскільки система спочатку орієнтована на математичну обробку інформації, істотно полегшуються операції з трансформації даних, а також з відображення результатів роботи моделі, включаючи тривимірні графіки.

Реалізація генетичних алгоритмів в складі MatLab дозволяє обмежитися тільки написанням функції пристосованості. При цьому хромосома задається як набір змінних одного з стандартних типів, що спрощує процес кодування / декодування. Інші параметри алгоритму можна задати в інтерактивному режимі, як і в Sugal. Кількість різних параметрів настройки також велика.

Все це робить MatLab простішим в освоєнні, універсальнішим і зручнішим у використанні, ніж Sugal. Разом з тим, процес моделювання в системі MatLab все ще вимагає наявності навичок програмування, хоча і в меншій мірі, ніж для низькорівневого програмного забезпечення [148].

Пакет Evolver є представником вбудованого програмного забезпечення реалізації генетичних алгоритмів. Модель вводиться у вигляді пов'язаних між собою цільової функції, змінних і обмежень на стандартному аркуші Microsoft Excel. Управління роботою Evolver проводиться через панель інструментів, звідки можна задати параметри моделі, визначити параметри оптимізації, складання звітів і власне генетичного алгоритму [190].

Робота з Evolver схожа зі складанням моделі лінійного програмування, а сама оптимізація відбувається прозоро для користувача що істотно спрощує користування даним програмним продуктом. Разом з тим для моделей в Evolver діють всі основні переваги генетичних алгоритмів, наприклад, довільний вид функцій, що оптимізуються. Таким чином даний програмний продукт є зручним для використання в тих випадках, коли для вирішення задачі необхідно використовувати можливості генетичних алгоритмів, але немає сенсу втрачатись у тонкощі програмування.

Практичний інтерес представляє зіставлення швидкості роботи розглянутих програм. Відповідно до розглянутих в п. 2.3 підходів, тестування швидкості і якості роботи програмних продуктів необхідно проводити при вирішенні типових задач. Оскільки однією з основних переваг генетичних алгоритмів є ефективне (у порівнянні з іншими методами) рішення NP-повних задач, в якості типової задачі, що задовольняє вимогам В.2.1 – В.2.6 візьмемо задачу комівояжера.

Ця задача включена в приклади, що додаються до всіх розглянутих програмних пакетів. При цьому, однак, можна відзначити, що в Sugal і Matlab розташування міст задається координатами, що більш зручно, ніж в Evolver, де у вхідній таблиці необхідно задати відстані між усіма населеними пунктами. Крім того, Sugal і Matlab за результатами оптимізації виводять на екран карту, на якій показано кращий зі знайдених алгоритмом маршрутів.

Час роботи алгоритмів перевірявся при вирішенні задачі комівояжера на карті з 20 міст при 1000 поколінь та популяції з 50 особин. Вирішення цієї задачі методом перебору вимагає розгляду $6 \cdot 10^{16}$ комбінацій маршрутів, однак для генетичних алгоритмів збіжність настає настільки швидко, що точний вимір часу роботи програми на сучасній ЕОМ стає неможливим. Тому тестування проводилося на застарілому комп'ютері з процесором Pentium - II, що працює на частоті 366 МГц. Його результати показано в табл. 4.6.

Таблиця 4.6 – Час роботи генетичних алгоритмів при вирішенні задачі комівояжера

№пп	Програмний пакет	Час роботи, сек.	Примітка
1.	Sugal 2.1	4	–
2.	Matlab 7.0	26	Алгоритм зійшовся вже через 200 поколінь
3.	Evolver 5.0	35	Для сходження алгоритму знадобилося 3000 поколінь

Як видно з табл. 4.6, найкращі за часом результати показав алгоритм, реалізований на базі пакету Sugal. У всіх тестах, після тисячі поколінь алгоритм зійшовся, тобто значення функції пристосованості переставало поліпшуватися (рис. 4.2). Однак зазначимо, що оптимальне рішення алгоритм знаходив не завжди. Іноді обраний алгоритмом маршрут відрізнявся від найкоротшого (знайденого за результатами інших запусків) на 10-15%.

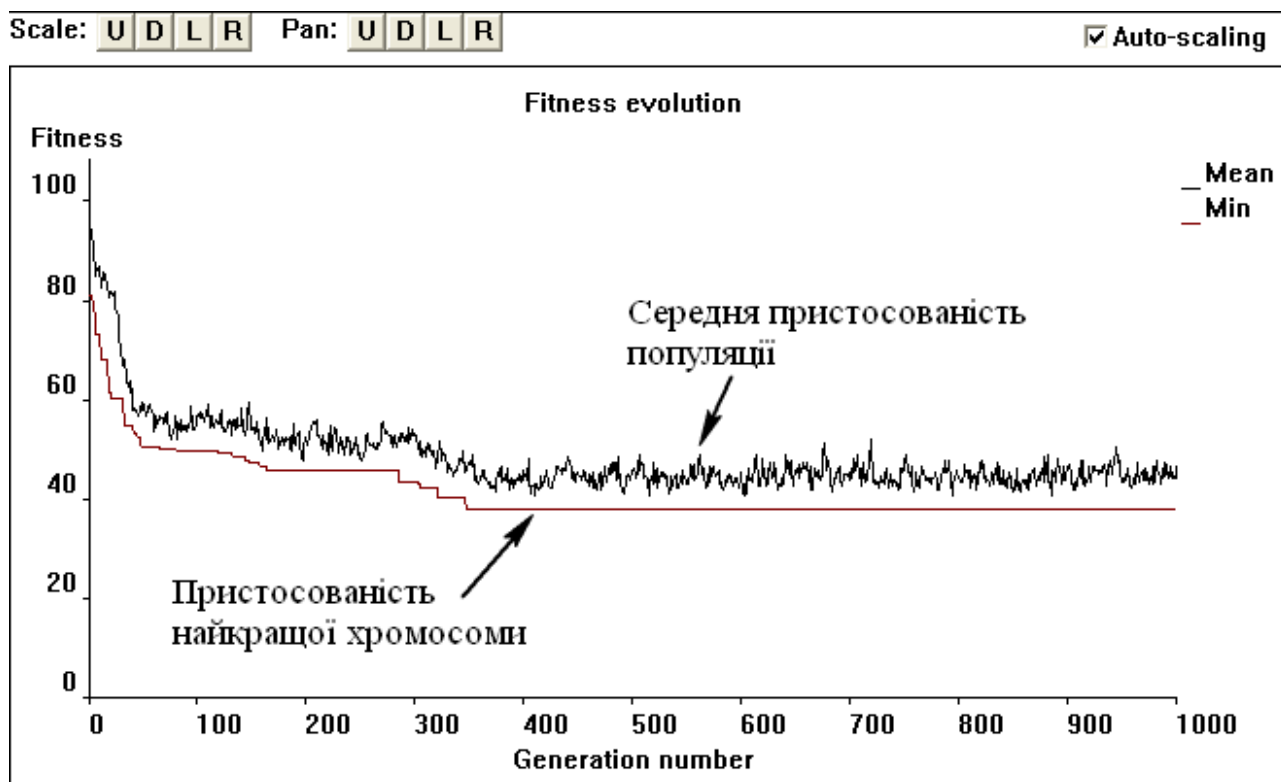


Рисунок 4.2 – Динаміка зміни пристосованості популяції при вирішенні задачі комівояжера в Sugal.

Алгоритм, реалізований в пакеті Matlab показав другий результат за часом роботи, значно поступившись за цим параметром Sugal. Однак якість роботи Matlab виявилось набагато вище. Так, аналіз динаміки значень функції пристосованості (рис. 4.3) показує, що алгоритм зійшовся вже через 200 поколінь, що є найкращим результатом серед розглянутих програм.

Суб'єктивний аналіз карти з маршрутом об'їзду міст дозволяє зробити висновок про те, що знайдене рішення є дуже близьким до оптимального.

Такий результат можна пояснити реалізацією в Matlab сучасних моделей кросоверу і селекції, що поліпшує збіжність алгоритму.

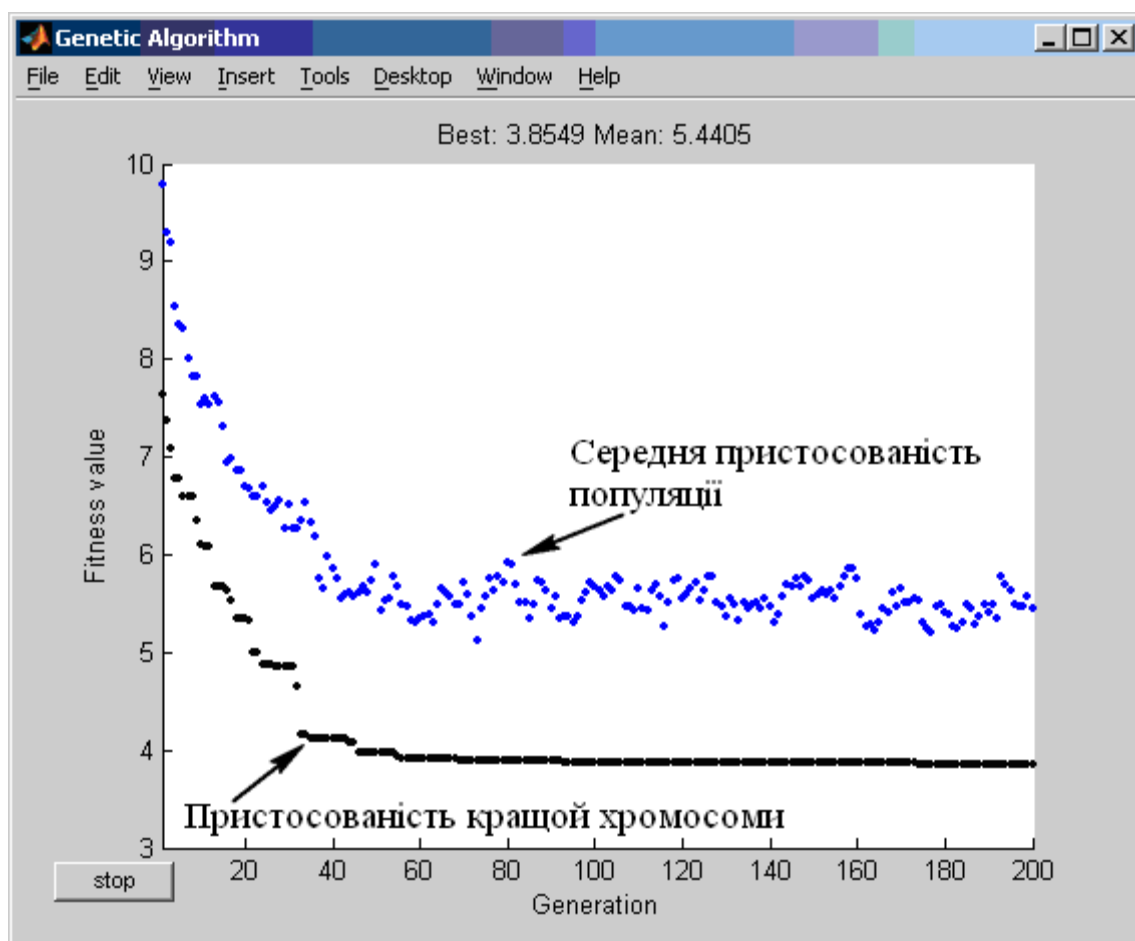


Рисунок 4.3 – Динаміка зміни функції пристосованості при вирішенні задачі комівояжера в Matlab

Однак реалізація генетичного алгоритму в Matlab має і недоліки. При включенні відображення проміжних результатів роботи (карти маршруту і значення функції пристосованості), швидкість розрахунків сповільнюється приблизно в 15 разів. Так, при інших рівних параметрах, на 200 поколінь було витрачено 78 секунд.

Найповільнішим виявився пакет Evolver. Незважаючи на те, що за заявою виробника даного ПЗ, швидкість його роботи в порівнянні з попередніми версіями зросла приблизно в 20 разів [51], за часом рішення тестової задачі Evolver відстає від Matlab. Швидкість збіжності алгоритму розв'язання задачі також виявилася не дуже хорошою, оскільки сходження наступало тільки після 3000 поколінь (рис. 4.4), тоді як в Sugal - після 500 поколінь, а в MatLab - після 200).

Низька швидкість роботи алгоритму обумовлена особливостями пакету Microsoft Office, що не призначений для таких вибагливих до ресурсів завдань, як генетична оптимізація. Найбільш імовірною причиною поганої збіжності, ймовірно, є застосування розробниками спрощених варіантів генетичних алгоритмів. Сучасні моделі генетичного пошуку (гібридні, багаторівневі,

паралельні і ін.), що поліпшують збіжність, але вимагають великих обчислювальних потужностей, не фігурують ні в описі Evolver, ні в списку можливих налаштувань.



Рисунок 4.4 – Динаміка зміни функції пристосованості при вирішенні задачі комівояжера в Evolver.

Порівняємо тепер вартість програмних продуктів генетичної оптимізації (табл. 4.7).

Таблиця 4.7 – Вартість програмних продуктів для розв'язання задач генетичної оптимізації

№ з.п.	Назва продукту	Ціна за версію, USD				Наявність інших версій
		Тільки генетичні алгоритми		Разом із нейронними мережами		
		Комерц.	Освітня.	Комерц.	Освітня	
1	Gene Hunter	595	—	1395	—	—
2	Evolver	£ 750 =\$968*	-50% =\$484	£1245=\$1606 £1850**=\$2386	-50%	пробна, студентська
3	MathLab + ГА ***	2350+1200 =3550	от \$55	2350+1200+1200 =4750	от \$55	—
4	Sugal	0	0	—	—	—

* Вартість одного фунта стерлінгів на момент проведення дослідження становила 1,29 долара США

** в складі повного пакета програм DecisionTools Suite, куди входить крім нейронних мереж і генетичних алгоритмів ще чотири програмних продукту - аналіз ризиків, дерева рішень, статистика і метод Монте-Карло.

*** складається з вартості базового пакета MatLab і вартості відповідних доповнень.

Крім розглянутих вище продуктів, в табл. 4.7 включений також пакет генетичної оптимізації *Gene Hunter*, що випускається компанією *Ward Systems* [191]. Як і *Evolver*, він відноситься до категорії вбудованого ПЗ і має схожий інтерфейс, однак компанія-виробник не пропонує ознайомчу версію програми, тому провести тестування швидкості її роботи не виявилось можливим. Крім того, оскільки практично всі виробники разом із реалізацією генетичних алгоритмів пропонують нейромережеві програмні продукти, зроблено аналіз їх загальної вартості.

Як видно з аналізу табл. 4.7, з позиції вартості, найбільші переваги має пакет *Sugal*, оскільки може бути використаний абсолютно безкоштовно. З інших програмних продуктів виділити однозначного лідера важко. При покупці програми тільки для задач генетичної оптимізації найнижча вартість у *GeneHunter*. При необхідності комплексного вирішення завдань штучного інтелекту, оптимальним вибором для навчальних закладів буде покупка *MatLab*. Для підприємств ж оптимальним рішенням і в цьому випадку залишається придбання *GeneHunter*. Однак слід враховувати, що до складу повного пакета *DecisionTools Suite* крім нейронних мереж і генетичних алгоритмів (*Evolver*) входять і інші методи аналізу даних, що в деяких випадках може вплинути на прийняття рішення про покупку програмного продукту.

Таким чином, можна систематизувати переваги і недоліки розглянутих типів програмного забезпечення для моделювання генетичних алгоритмів (табл. 4.8).

Таким чином, можна зробити наступні висновки про доцільність використання різних типів програмного забезпечення:

Низькорівневе ПЗ (Sugal) відрізняється високою швидкодією, і поширюється на вільній основі. Це робить його оптимальним при вирішенні окремих завдань, коли важлива низька вартість ПЗ, або при вирішенні безлічі однотипних завдань, що відрізняються лише вхідними даними, коли важливим виявляється висока швидкодія. Вирішальною умовою при виборі розглянутого пакета *Sugal* є можливість написання програм на мові C++.

Спеціальне ПЗ (MatLab) – є одним з найпотужніших сучасних математичних інструментів. Для роботи з *MatLab* досить елементарних навичок програмування на мові типу Basic. Трудомісткість роботи після освоєння пакету знаходиться на рівні вбудованого ПЗ. У той же час можливості пакету допускають рішення задач будь-якої складності. Таким чином, використання *MatLab* оптимально в навчальних та наукових організаціях, а також в установах, які займаються комерційним використанням генетичних алгоритмів для вирішення практичних завдань.

Вбудоване ПЗ (Evolver) є найбільш простим в освоєнні і використанні інструментом оптимізації, що заснована на використанні ГА, оскільки інтегрується в стандартний офісний програмний продукт *Microsoft Excel* і забезпечує користувачеві знайоме середовище розробки. Це єдиний програмний продукт з розглянутих, що дозволяє вирішувати задачі в режимі, близькому до інтерактивного. Перевагою *Evolver* є наявність пробної версії пакета, завдяки чому можна протягом 15 днів, без фінансових витрат вивчити

основні можливості оптимізації, що заснована на використанні ГА.

Таблиця 4.8 – Характеристика основних типів програмного забезпечення реалізації генетичних алгоритмів в ІСПР

№ _{пп}	Тип	Представ-ники	Переваги	Недоліки
1	Низькорівневе	Sugal	Найвища швидкість розрахунків. При наявності кваліфікованих програмістів може бути розширено новими функціями. Розповсюджується безкоштовно.	Потрібний кваліфікований програміст. Велика трудомісткість вирішення задач. Відсутня підтримка і оновлення ПЗ.
2	Спеціальне	MatLab	Швидка збіжність алгоритму. Регулярні оновлення, з урахуванням останніх наукових досягнень. У поєднанні з іншими компонентами пакету утворює найпотужніший комплекс математичного програмного забезпечення	Більш ніж середня трудомісткість освоєння продукту і рішення задач. Необхідні навички програмування. Висока вартість.
3	Вбудоване	Evolver, GeneHunter	Простота освоєння і вирішення завдань малої і середньої складності.	Низька швидкість. Повільна збіжність алгоритму. Не реалізовані сучасні досягнення в області ГА. Висока вартість.

У той же час розглянутому пакету властиві такі недоліки, як низька швидкість роботи і триваліша, ніж у інших розглянутих продуктів, збіжність алгоритму. Таким чином, областю застосування Evolver є задачі малого та середнього рівня складності. Його використання доцільне, якщо є необхідність в оптимізації з використанням генетичних алгоритмів, але немає можливості для детального вивчення всіх тонкощів їх застосування. Evolver також підходить для попереднього моделювання економічних задач в ІСПР.

Отже сучасне програмне забезпечення інтелектуальних обчислень дозволяє забезпечити реалізацію модулів ІСПР із використанням готових бібліотек та компонентів, що знижує витрати часу та трудомісткість розробки ІСПР та її інтеграції із інформаційними системами суб'єктів економічних відносин.

4.2 Інтелектуальні методи прогнозування показників бюджетного процесу

Бюджетно-податкова система України здійснює перерозподіл величезної маси грошових коштів, які держава стягує у вигляді податків з підприємств та громадян і потім повертає в економіку через механізм державних витрат та видатків. Процес отримання податків і здійснення видатків, передбачених у затвердженому бюджеті має назву «виконання бюджету». Зазвичай під час виконання бюджету спостерігається відхилення від затвердженого варіанта, у зв'язку з чим необхідно контролювати та регулювати цей процес [194, с.50-51]. Виконати бюджет – означає точно у визначені строки та в установлених розмірах мобілізувати заплановані доходи і здійснити передбачені видатки. Але, внаслідок дій різноманітних факторів, фактичні показники виконання бюджету майже завжди відрізняються від запланованих. Деякі з цих факторів пов'язані із непередбачуваними обставинами, але значну частину з них можна виявити на стадії складання бюджету і їх прояв свідчить про недосконалість бюджетного процесу.

Головним бюджетом країни є державний бюджет, який через систему причинно-наслідкових зв'язків має вплив на всі сфери народного господарства. Тому для державного бюджету питання реалістичності цільових показників набуває найбільшого значення. Але аналіз фактичних даних про виконання державного бюджету з 2000 року (табл. 4.9) свідчить про те, що відхилення від планових результатів іноді досягало 18% (у 2009 році), а відхилення у межах 1% спостерігалось лише у 6 випадках з 32 [106].

Таблиця 4.9 – Результати виконання плану доходів та видатків державного бюджету України за 2000-2015 роки

Рік	Факт. викон. доходів бюджету	Факт. викон. видатків бюджету	Рік	Факт. викон. доходів бюджету	Факт. викон. видатків бюджету
2000	106,7	104,7	2008	97,1	91,7
2001	92,8	95,9	2009	82,1	85,2
2002	100,2	111,6	2010	96,5	93,5
2003	103,4	100,4	2011	100,4	94,1
2004	107,9	97,3	2012	90,3	92,7
2005	97,5	94,6	2013	96,7	99,3
2006	100,7	93,4	2014	94,6	93,3
2007	98,6	93,6	2015	100,5	96,2

Взагалі ж дані табл. 4.9 свідчать про те, що у держави немає інструменту оцінювання реалістичності плану доходів та видатків державного бюджету.

Таке важливе питання, як прогнозування рівня виконання бюджету, зрозуміло знайшло своє відображення в працях вітчизняних та закордонних науковців. Завдяки їх працям було розвинуто методи індуктивного прогнозування виконання бюджетних показників, тобто методи «від часткового

до загального».

З погляду економічної теорії загальні доходи та видатки бюджету є одним з макроекономічних показників, а згодом між ним, та іншими макроекономічними показниками повинен бути зв'язок. Серед чинників, які мають вплив на виконання державного бюджету чимало макроекономічних факторів. Але й досі в вітчизняних літературних джерелах відсутній математичний аналіз цих зв'язків та методи прогнозування рівня виконання держбюджету на їх основі [171].

Таким чином, актуальним є пошук методів оцінювання реалістичності планових показників бюджетно-податкової системи держави на підставі макроекономічних показників розвитку економіки. Для цього передбачається [171]:

- відібрати низку факторів, які теоретично можуть мати вплив на рівень виконання держбюджету;
- проаналізувати фактичний рівень впливу цих факторів;
- побудувати моделі їх впливу на фактичний рівень виконання держбюджету;
- оцінити якість моделей, які побудовано;
- зробити прогноз виконання держбюджету та перевірити його достовірність.

Аналіз підходів до бюджетного планування [229, 214, 202, 234] дозволяє виділити низку факторів, які впливають на бюджетний процес. У першому приближенні ці фактори можна поділити на дві групи: фактори, які обчислюються та фактори, які не обчислюються.

До факторів, які обчислюються, зокрема, відносяться [171]:

- рівень ВВП;
- індекс споживчих цін;
- індекс цін виробників промислової продукції;
- індекс реальної заробітної плати;
- рівень безробіття; сальдо торговельного балансу;
- розміри експорту та імпорту товарів і послуг;
- відносний рівень продукції промисловості; оборот роздрібної торгівлі.

До факторів, які не обчислюються відносяться [171]:

- політичні фактори;
- фінансове та бюджетне законодавство;
- податкове законодавство;
- митне законодавство;
- кліматичні умови;
- соціальна стабільність;
- кваліфікація і досвід керівників та фахівців, що розробляють бюджет;
- освіта й система перепідготовки кадрів;
- система управління якістю;
- забезпеченість кадрами;
- форс-мажорні обставини.

Для розробки ІСПР найбільше значення мають фактори, які

обчислюються, тому що їх аналіз може бути формалізований із допомогою методів економіко-математичного моделювання, а отримані результати використані для побудови прогнозу на наступні періоди. Тому далі будемо розглядати лише цю групу факторів, як найбільш прийнятну при розробці ІСПР.

Інформаційною базою дослідження є відкриті дані звітності НБУ [95], Держкомстату України [106] та інших джерел [111]. Оптимальним періодом для початку вибірки даних обрано 2000 рік. До цього розвиток економіки мав мінливий стан, тому дані можуть відображати некоректну інформацію. До того ж облік деяких макроекономічних показників до цього періоду не здійснювався. На підставі зазначених даних складено базу макроекономічних показників для наступного аналізу (табл. 4.10).

В табл. 4.10. містяться дані, які мають різну природу. Так показники фактичного виконання доходів та видатків бюджету, рівня безробіття та сальдо торгівельного балансу наведено за їх поточними значеннями, а показники валового внутрішнього продукту, індексу споживчих цін, цін виробників промислової продукції, індексу реальної заробітної плати, експорту і імпорту товарів і послуг, приросту промислової продукції, обігу роздрібною торгівлі та рівня інфляції відображенні темпами приросту в порівнянні з минулим роком.

Таблиця 4.10 – Основні макроекономічні показники розвитку економіки України в 2000-2014 р.

Рік	Факт. викон. доходів бюджету	Факт. викон. видатків бюджету	ВВП реальн.	Індекс споживчих цін	Індекс цін виробників промислової прод.	Індекс реальної ЗП
2000	106,7	104,7	105,9	125,8	120,8	99,1
2001	92,8	95,9	109,2	106,1	100,9	119,3
2002	100,2	111,6	105,2	99,4	105,7	118,2
2003	103,4	100,4	109,6	108,2	111,1	115,2
2004	107,9	97,3	112,1	106,3	115,3	123,8
2005	97,5	94,6	102,6	110,3	109,5	120,3
2006	100,7	93,4	107,1	111,6	114,1	118,3
2007	98,6	93,6	107,6	116,6	123,3	112,5
2008	97,1	91,7	102,1	122,3	123	106,3
2009	82,1	85,2	85,2	112,3	114,3	100,9
2010	96,5	93,5	104,2	109,1	118,7	110,2
2011	100,4	94,1	105,2	104,6	114,2	108,7
2012	90,3	92,7	100,2	99,8	100,3	114,4
2013	96,7	99,3	100	104,8	105,5	103,3
2014	94,6	93,3	93,2	124,9	131,8	93,5

Продовження табл. 4.10

Рік	Рівень безробіття	Сальдо торговельного балансу, млн. USD	Експорт товарів і послуг	Імпорт товарів і послуг	Індекс пром. прод.	Оборот роздрібн. торг.	Рівень інфляції
2000	12,4	2742,5	118,8	118,2	112	107	125,8
2001	11,7	661,5	82,8	93,3	114	115	106,1
2002	10,3	1039,8	108	105,7	107	117	99,4
2003	9,7	335,1	128,1	134,7	116	121	108,2
2004	9,2	3412,4	142,6	128,1	112	120	112,3
2005	7,8	-1340,2	105,6	124,6	103	122	110,3
2006	7,4	-5222,9	111,4	122,8	106	127	111,6
2007	6,9	-9592	128,2	135,4	110	129	116,6
2008	6,9	-17710,9	141,8	149,4	95	118	122,3
2009	9,6	-4815	56,5	50	78	83	112,3
2010	8,8	-7958	129,6	133,8	111	110	109,1
2011	8,6	-12764,2	134,3	138,1	108	115	104,6
2012	8,1	-13776	101,8	102,9	99,5	115	99,8
2013	7,7	-12606,5	91	91,1	96	110	100,5
2014	9,7	-5283	83	66,2	91,5	101,4	124,9

Для економіко-математичного аналізу та моделювання впливу зазначених чинників на процес виконання державного бюджету дані табл. 4.10 слід додатково обробити. По-перше їх слід розширити за рахунок введення відносних та прирістних показників. По-друге, оскільки дані будуть використані для прогнозування слід виконати їх зсув по рокам. Необхідність зсуву даних ґрунтується на тому, що на початок кожного поточного року в наявності є тільки дані по макроекономічних показниках минулого року та план бюджету на поточний рік.

В табл. 4.10. містяться дані, які мають різну природу. Так показники фактичного виконання доходів та видатків бюджету, рівня безробіття та сальдо торговельного балансу наведено по їх поточним значенням, а показники валового внутрішнього продукту, індексу споживчих цін, цін виробників промислової продукції, індексу реальної заробітної плати, експорту і імпорту товарів і послуг, приросту промислової продукції, обігу роздрібно торгівлі та рівня інфляції відображенні темпами приросту в порівнянні з минулим роком.

Для економіко-математичного аналізу та моделювання впливу зазначених чинників на процес виконання державного бюджету дані табл. 4.10 слід додатково обробити. По-перше їх слід розширити за рахунок введення відносних та прирістних показників. По-друге, оскільки дані буде використано

для прогнозування слід виконати їх зсув по рокам. Необхідність зсуву даних ґрунтується на тому, що на початок кожного поточного року в наявності є дані по макроекономічних показниках минулого року та план бюджету на поточний рік.

На всіх етапах аналізу використовується програма Deductor Studio Academic, яка має широкий вибір вбудованих методів аналізу, та дозволяє безкоштовне використання.

В процесі аналізу першим чином виявляється взаємозв'язок між макроекономічними показниками та їх впливом на виконання державного бюджету. Результати кореляційного аналізу даних табл. 4.9 у вигляді, що надає пакет Deductor Studio Academic, наведено на рис. 4.5.

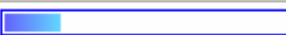
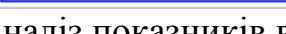
Входные поля		Корреляция с выходными полями	
№	Поле	Доходи ф/пл	Видатки ф/пл
1	д.п./ВВП	 -0,194	 -0,553
2	в.п./ВВП	 -0,178	 -0,657
3	ВВП ном.отн.изм	 -0,142	 -0,171
4	експ/ВВП	 0,246	 0,300
5	имп/ВВП	 -0,269	 -0,143
6	ВВП реальний	 0,148	 0,148
7	індекс споживчих цін	 -0,309	 -0,368
8	індекс цін виробн. пром. прод.	 -0,359	 -0,659
9	індекс реальної ЗП	 0,228	 0,177
10	рівень безробіття	 0,411	 0,650
11	сальдо/ВВП	 0,496	 0,488
12	експорт товарів і послуг	 -0,286	 -0,485
13	імпорт товарів і послуг	 -0,296	 -0,478
14	пром.прод.	 0,353	 0,313
15	оборот роздрібн.торг.	 -0,042	 -0,157
16	рівень інфляції	 -0,282	 -0,351

Рисунок 4.5 – Кореляційний аналіз показників виконання державного бюджету

Занадто велика кількість вхідних показників у поєднанні із короткою вибіркою даних може ускладнити моделювання. Тому за результатами кореляційного аналізу потрібно зробити її проріджування: відкинути ті дані, які мають східне походження, або незначний зв'язок із результуючими показниками. Так, дані по експорту та імпорту товарів з погляду виконання бюджету найкраще відображає показник сальдо експорту/імпорту, що дозволяє залишити один показник з п'яти присутніх у вхідних даних. Аналогічним чином до скороченої вибірки даних додамо показник рівня безробіття, індекс цін виробників промислової продукції, показники планових доходів та видатків бюджету, а також показник динаміки реального ВВП.

Розглянемо регресійну модель з урахуванням зазначених макроекономічних показників. Із допомогою пакету Deductor Studio Academic було отримано дві такі моделі: перша по доходах, друга – по видатках.

Результати моделювання наведено на рис.4.6.

Выходное поле: Доходы ф/пл	
Атрибут	Коэффициент
9.0 <Константа>	100,95
9.0 д.п./ВВП	-28,507
9.0 в.п./ВВП	97,769
9.0 ВВП реальный	0,15703
9.0 индекс цін виробн. пром. прод.	-0,44937
9.0 рівень безробіття	1,1091
9.0 сальдо/ВВП	53,754

Выходное поле: Видатки ф/пл	
Атрибут	Коэффициент
9.0 <Константа>	108,13
9.0 д.п./ВВП	105,07
9.0 в.п./ВВП	-45,227
9.0 ВВП реальный	0,090911
9.0 индекс цін виробн. пром. прод.	-0,56802
9.0 рівень безробіття	2,9554
9.0 сальдо/ВВП	-24,993

Рисунок 4.6 – Коэффициенты регрессии моделей выполнения бюджета по доходам та видаткам

Таким чином, за результатами моделювання, лінійна регресійна модель виконання плану доходів має наступний вигляд:

$$\text{Виконання доходу} = 100,95 - 28,507 \cdot x_1 + 97,769 \cdot x_2 + 0,15703 \cdot x_3 - 0,44937 \cdot x_4 + 1,1091 \cdot x_5 + 53,754 \cdot x_6, \quad (4.1)$$

де x_1 – д.п./ВВП; x_4 – індекс цін виробників промислової продукції;
 x_2 – в.п./ВВП; x_5 – рівень безробіття;
 x_3 – ВВП реальний; x_6 – сальдо/ВВП.

Лінійна регресійна модель виконання плану видатків має вигляд:

$$\text{Виконання видатків} = 108,13 + 105,07 \cdot x_1 - 45,227 \cdot x_2 + 0,090911 \cdot x_3 - 0,56802 \cdot x_4 + 2,9554 \cdot x_5 - 24,993 \cdot x_6. \quad (4.2)$$

Для аналізу якості моделі побудуємо діаграму розсіювання. Для моделей лінійної регресії діаграми розсіювання наведено на рис. 4.7.

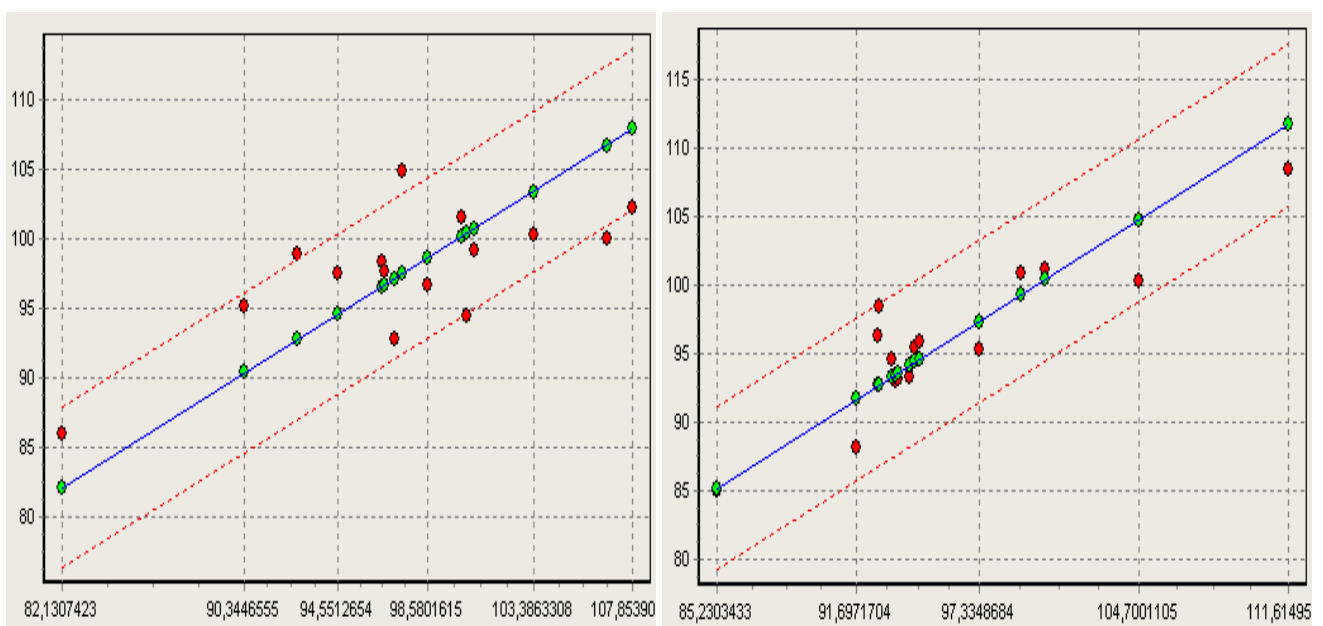


Рисунок 4.7 – Діаграми розсіювання моделі лінійної регресії по доходах (ліворуч) та по видатках (праворуч)

Як випливає з аналізу діаграм розсіювання, моделі (4.1) та (4.2) дозволяють отримати прогнозні показники, які в цілому укладаються в межі довірчого інтервалу 5% (пунктирна лінія). Також, можна бачити, що модель лінійної регресії по видатках має більшу адекватність, ніж модель по доходах.

Істотним недоліком моделей лінійної регресії є принципова неможливість урахування складних залежностей між даними, які у реальній економіці можуть бути далекими від лінійних. Тому далі побудуємо та проаналізуємо нейромережеві моделі, які привабливі тим, що можуть моделювати майже усі види залежностей між даними.

Оскільки одним з недоліків використання нейронних мереж для прогнозування є неможливість заздалегідь визначити, яка конфігурація мережі визначиться найкращою, необхідно побудувати й проаналізувати декілька нейронних мереж та обрати ту, яка дає кращий результат за помилкою навчання та діаграмою розсіювання.

На рис. 4.8 наведено структуру нейронної мережі, яка навчалась на тих самих даних, що моделі лінійної регресії (4.1) та (4.2).

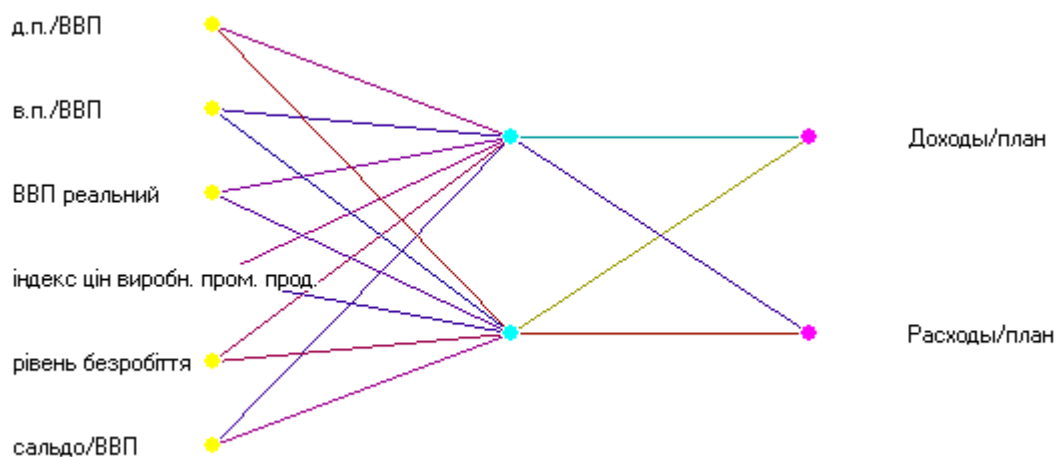


Рисунок 4.8 – Нейронна мережа із шістьма показниками

Дана модель (рис. 4.8) має класичний вигляд, та формулу 6-2-2. Середня помилка прогнозування для цієї моделі складає 0,00642. Діаграми розсіювання по доходах та видатках наведено на рис.4.9

Порівняння діаграм (рис. 4.9) із діаграмами (рис. 4.7) показує, що шестифакторна нейромережева модель прогнозування виявилась більш точною, ніж шестифакторні лінійно-регресивні моделі (4.1) та (4.2). Більшість показників на рис. 4.9 розташовано біля лінії ідеальних значень.

Для того, щоб ще підвищити точність прогнозування побудуємо декілька нейромережевих моделей із різною структурою та різними наборами вхідних показників. Результати моделювання за ознакою середньої помилки навчання наведено у табл. 4.11.

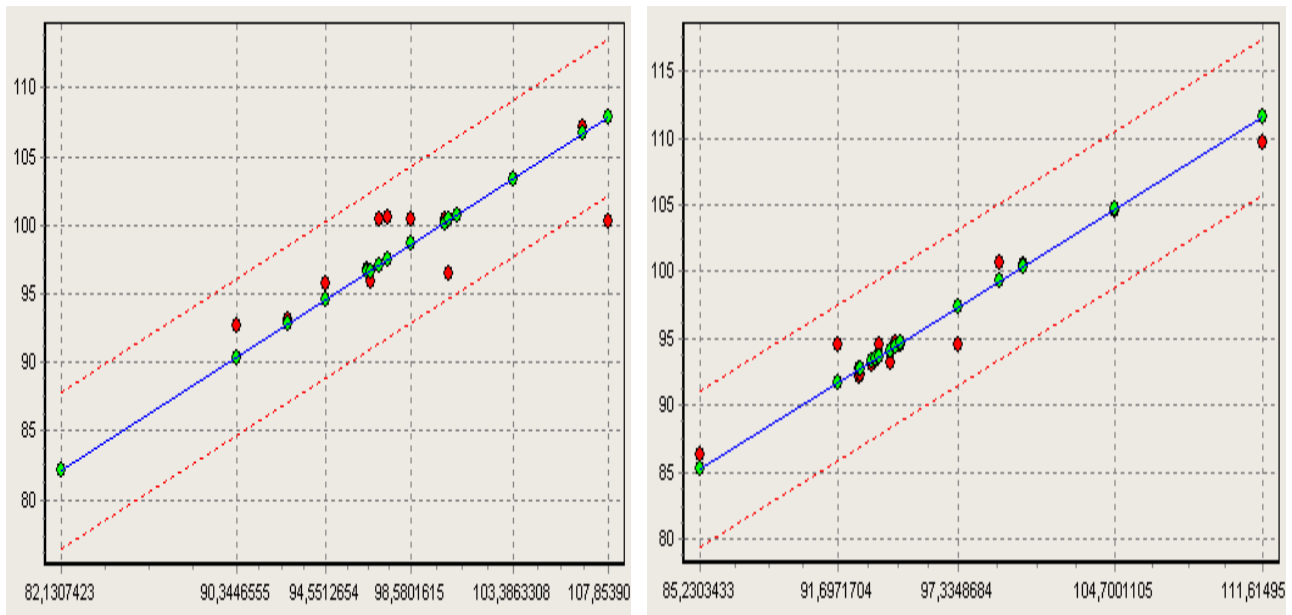


Рисунок 4.9 – Діаграми розсіювання шестифакторної нейромережевої моделі по доходах (ліворуч) та по видатках (праворуч)

Таблиця 4.11 – Результати навчання нейронних мереж у різній конфігурації

№пп	Кількість входних показників	Кількість нейронів в прихованому шарі	Середня помилка навчання.
1	6	2	0,00642
2	16	2	0,000303
3	16	3	0,0000858
4	10	2	0,00345
5	10	3	0,0000376

Під час моделювання досліджувалися й інші конфігурації, але для уникнення надмірного збільшення обсягу роботи, обрано лише типові з них.

Таким чином, кращою, за критерієм середньої помилки виявилась нейронна мережа із 10 показниками, що аналізуються, та 3 нейронами у прихованому шарі, формула якої 10-3-2. Структуру такої мережі показано на рис. 4.10, а відповідні діаграми розсіювання на рис. 4.11.

Порівняння діаграм (рис. 4.11) із наведеними вище (рис. 4.9), свідчить про те, що дана нейромережева модель діє найкращим чином, ніж попередні. Вона добре описує стан економіки України, та показує найдостовірніші дані про виконання державного бюджету в Україні за останні роки. Головним недоліком десятифакторної моделі є потенційна схильність до перенавчання, що впливає з зіставлення розміру навчальної вибірки із результатами розрахунків за формулами (2.5) і (2.25).

Для побудови прогнозу виконання державного бюджету України будемо використовувати розглянуту десятифакторну нейромережеву модель. Оскільки для навчання використовувалися дані за 2000-2014 роки, спрогнозуємо показники виконання бюджету на 2015 рік. Вхідні дані наведено у табл. 4.12.

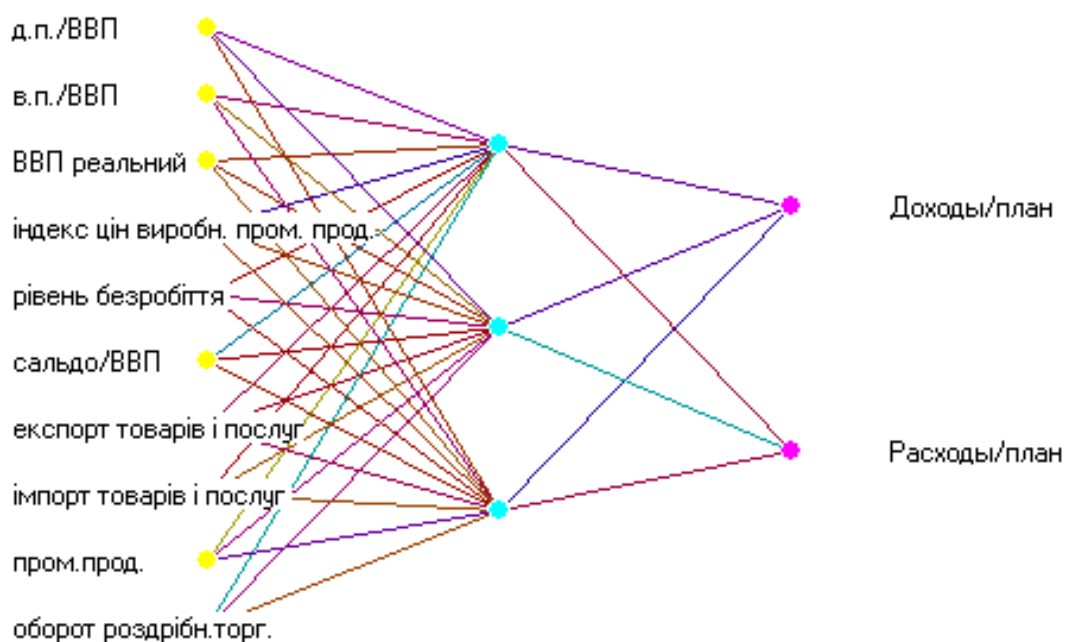


Рисунок 4.10 – Структура нейронної мережі із 10 показниками

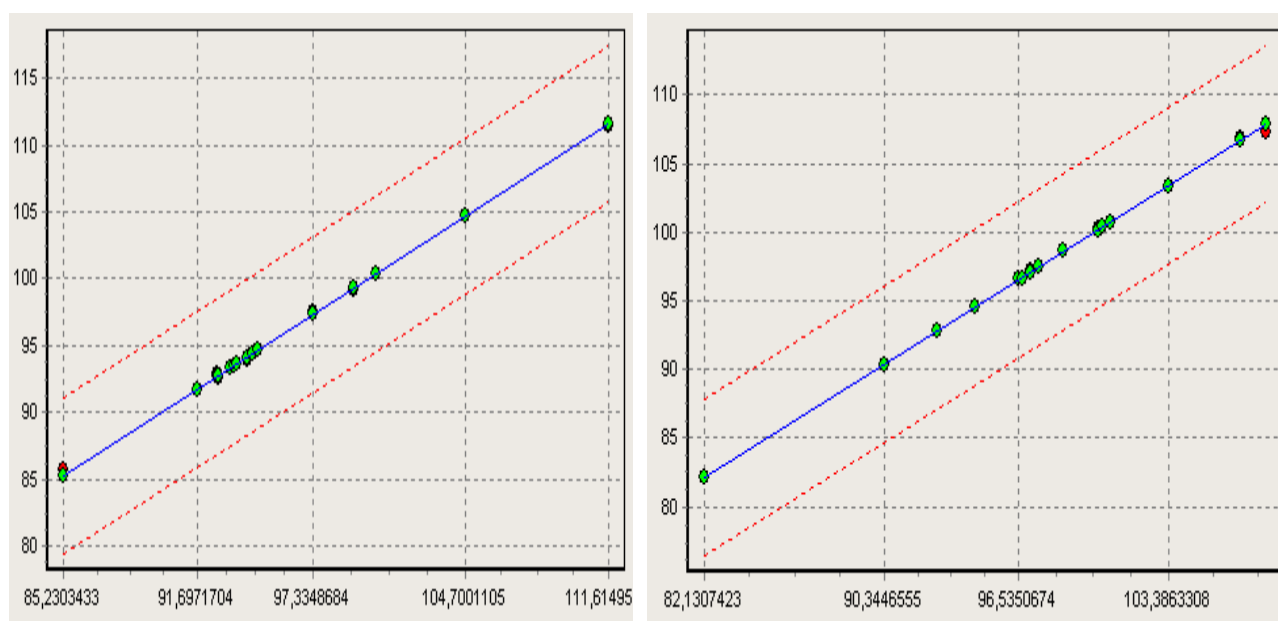


Рисунок 4.11 – Діаграма розсіювання при прогнозуванні видатків (ліворуч) та доходів (праворуч) для нейронної мережі із 10 показниками

Таблиця 4.12 – Вхідні дані для прогнозування показників бюджетно-податкової системи (станом на початок 2015 року)

Показник	Значення	Показник	Значення
Доходи планові, млн. грн.	504993,3	Рівень безробіття, %	9,7
Видатки планові, млн. грн.	569768	Сальдо торг. балансу млн. USD	-5283
ВВП номінальний, млн. грн.	1566728	Експорт товарів і послуг, %	83
Експорт, млн.USD	47695,3	Імпорт товарів і послуг, %	66,2
Імпорт млн.USD	46389,8	Випуск промислової продукції	91,5
ВВП реальний, %	93,2	Оборот роздрібної торгівлі	101,4
Індекс цін виробн. пром. прод., %	131,8		

На підставі цих даних, які внесені до моделі із допомогою інструменту «что-если» системи Deductor Studio Academic, отримано прогноз, який подано на рис 4.12.

Поле	Значение
Входные	
9.0 д.п./ВВП	0,32232353
9.0 в.п./ВВП	0,363667465
9.0 ВВП реальный	93,2
9.0 индекс цін виробн...	131,8
9.0 рівень безробіття	9,7
9.0 сальдо/ВВП	-0,047207939
9.0 экспорт товаров і...	83
9.0 импорт товаров і п...	66,2
9.0 пром.прод.	91,5
9.0 оборот роздрібн....	101,4
Выходные	
9.0 Доходи ф/пл	102,352765339118
9.0 Видатки ф/пл	96,8318972475095

Рисунок 4.12 – Прогнозні дані нейромережевої моделі

З рис. 4.12 можна побачити, що за прогнозом виконання державного бюджету в 2015 році складатиме 102.35% в доходній частині та 96.83% в частині видатків.

Порівняємо результати прогнозу (рис. 4.13) із реальними даними виконання державного бюджету. Так, офіційні дані про виконання державного бюджету України за 2015 рік наступні [98]:

1. За дванадцять місяців 2015 року доходну частину державного бюджету загалом виконано на 100,5% річного плану. При цьому слід зазначити, що впродовж 2015 року, суму доходів було збільшено з 475.9 млрд. грн. до 517 млрд. грн. Якщо взяти первинну суму запланованих доходів, то виконання цієї частини бюджету складатиме 108,6%.

2. Видаткову частину виконано на 96,2% річного плану.

Таким чином можна визначити, що в цілому прогнозні показники, одержані за допомогою нейромережевої моделі (рис. 4.10) практично співпадають с реальними даними [98]. Так, модель прогнозування, яку зроблено у роботі спрогнозувала, що план доходів бюджету буде перевиконано на 2.35%, а план видатків недовиконано на 3.2%. Фактичні дані свідчать про перевиконання плану доходів держбюджету на 0,5% (або 8,6% без урахування змін в бюджеті, впродовж року), та недовиконання плану видатків по загальному фонду на 3,8%. Це достатньо висока точність для моделі, із такою невеличкою кількістю показників, що аналізуються.

Також слід зазначити, що абсолютно точний прогноз не є можливим, оскільки на виконання бюджету окрім проаналізованих факторів впливає велика кількість показників, що не є обчислювальними, зокрема політичні

фактори, законодавство, кліматичні умови, соціальна стабільність, кваліфікація фахівців, що розробляють бюджет та інші.

Результати прогнозування показників виконання бюджетного процесу можуть бути використані виконавчими і законодавчими органами влади для прийняття рішень при формуванні і здійсненні фінансової політики, виявлення резервів для залучення коштів у бюджет, підвищення ефективності їх використання, посилення контролю за їх формуванням і витрачанням.

4.3 Генетичні методи оптимізації рефлексивних впливів промислових підприємств

Формальний підхід до розуміння і прогнозування процесів прийняття рішень було закладено із розробкою теорії ігор Дж. Фон Нейманом і О. Моргенштерном в 1944 році [178]. А уже в 1949 році сферу застосування теорії ігор істотно розширює Джон. Неш, який ввів поняття некооперативної рівноваги і розробив методи її аналізу [48]. Тим самим було доведено неоптимальність класичного «егоїстичного» підходу А. Сміта до ринкової конкуренції і почалося формування математичних інструментів економічного моделювання нового покоління. Подальший розвиток ігрового підходу до аналізу процесів прийняття рішень було продовжено зокрема розробкою принципів і методів рефлексивного управління [132], при якому враховуються не тільки об'єктивні фактори, а й схильності суб'єктів гри та їх бачення ситуації.

Рефлексивне управління отримало популярність як самостійна наукова дисципліна в зв'язку з роботами радянсько-американського вченого В.А.Лефевра. Розроблена ним аксіоматика і символічна система значно спрощує аналіз конфліктних взаємодій між суб'єктами у випадках, коли відомі такі їх характеристики, як інтенція, тиск зовнішнього середовища і т.п.

На підставі аналізу та прогнозування найбільш ймовірних дій противника, Лефевр запропонував впливати на мотиви його поведінки з тим, щоб перевести ситуацію в більш сприятливе для суб'єкта, що управляє, стан [132]. У його базових роботах рефлексивне управління трактується, як передача підстав противнику для прийняття ним вигідних для нас і шкідливих для нього рішень. Пізніше це визначення було розширено і, наприклад, в роботі [88] автори пропонують розуміти рефлексивне управління як специфічне інформаційно-психологічне управління зародженням і перетворенням смислів, рішень, інтенцій, цілей, цінностей, образів мислення і психологічних станів противника.

Зручне та ефективне використання рефлексивного управління для таких сфер, як міждержавні взаємини, зовнішня і внутрішня політика, військова справа неодноразово доводилася практично [220,133]. Використання ж даного інструментарію для аналізу та управління економічними взаєминами між підприємствами утруднено через низку обставин [147]:

- необхідно постійно враховувати дії безлічі агентів, які роблять свій вибір виходячи з різних передумов;

- обмеженість інформації про мотиви дій агентів і недостатня ефективність методів її отримання;
- необхідність враховувати економічну доцільність впливів і витрати на їх підготовку і забезпечення;
- велика кількість можливих каналів рефлексивних впливів за умов того, що доступні методи управління обмежені інформаційною рефлексією.

Але найсильнішою відмінністю економічних «битв» від військово-політичних є набагато слабша поляризація вибору у об'єкта рефлексивного управління. Іншими словами – поділ на «погане» і «добре», «моральне» і «аморальне» в економіці є дуже нечітким, а часто-густо просто неможливим. Це призводить до ускладнення процедури рефлексивного управління, за рахунок додавання до неї етапу поляризації альтернатив (рис. 4.13).

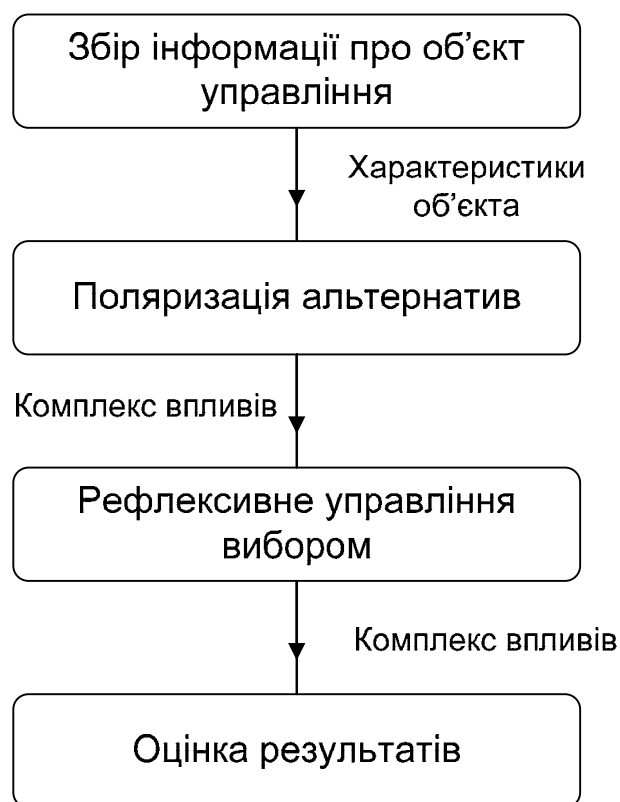


Рисунок 4.13 – Етапи рефлексивного управління в економічних системах

У схемі, яку представлено на рис. 4.13, рефлексивні дії розділені на два етапи. Перший етап є підготовчим і передбачає створення у контрагента штучної поляризації альтернатив, за якими йому згодом доведеться зробити вибір. При цьому альтернативам, які вигідні керуючому суб'єкту, присвоюється статус «добрих», а невигідним – статус «поганих». Класичним прикладом є кампанія «Підтримай вітчизняних товаровиробників», спрямована на створення позитивного образу продукції українських підприємств, незалежно від її якості, і споживчих властивостей. Якщо цілі етапу поляризації альтернатив успішно досягаються, то при здійсненні рефлексивного управління досить акцентувати увагу суб'єкта управління на характеристиках, які є «добрими» для керуючого суб'єкта.

Таким чином, при практичному використанні рефлексивного управління в економіці одним із важливих завдань є створення у суб'єкта управління такого *образу* простору вибору, в якому вигідна для керуючого суб'єкта альтернатива має найбільшу кількість переваг перед іншими.

Існуючий інструментарій рефлексивного управління [132] дозволяє лише проаналізувати можливі наслідки тих чи інших дій, та й то з деякими обмеженнями. Інакше кажучи, дослідник має можливість визначити результати своїх дій, але не має можливості визначити дії, які б привели до очікуваних результатів. Крім того, складність задачі зростає пропорційно кількості конкуруючих альтернатив.

Так, залежність кількості зв'язків (l), які необхідно аналізувати, від кількості контрагентів (n) визначається як:

$$l = \sum_{i=1}^n i, \quad (4.3)$$

або за формулою арифметичної прогресії

$$l = \frac{n + n^2}{2}. \quad (4.4)$$

Так, при одному або двох контрагентах (класичні випадки застосування рефлексивного управління, розглянуті в [132]) розглядаються відповідно один або три зв'язка. Але вже при десяти контрагентах кількість зв'язків зростає до 55. При цьому множина характеристик, відповідних кожному з зв'язків може мати досить велику розмірність, що практично виключає можливість аналізу без залучення програмно-технічних засобів.

Розглянемо постановку задачі управління образом простору вибору.

Нехай *Керуючим суб'єктом* є *Підприємство* - виробник деякої продукції (S). *Суб'єктом управління* є інше підприємство – потенційний *Споживач* даної продукції (B). *Конкурентами* (C) є підприємства – виробники аналогічної продукції. При цьому, оскільки на процедуру вибору впливають не тільки характеристики продукції, але і супутні обставини (терміни поставок, умови супроводу, гарантії, фінансові умови), доцільно розглядати не продукцію, а *Комерційні пропозиції* (КП) в цілому. Сукупність комерційних пропозицій Керуючого суб'єкта і Конкурентів становить простір вибору Суб'єкта управління.

Нехай Керуючому суб'єкту відомі Конкуренти, які беруть участь в боротьбі за виконання замовлення Суб'єкта управління, та характеристики їх комерційних пропозицій. Оскільки набори цих характеристик в більшості випадків не збігаються, Суб'єкт управління змушений відмовитися від їх безпосереднього зіставлення. Замість цього здійснюється угруповання характеристик і подальший зіставлення отриманих груп. Очевидно, що існує ціла позитивна кількість способів угруповання n , причому $n > 1$. Можна

припустити, що результат подальших дій в порівнянні груп, тобто оцінка Суб'єктом управління комерційних пропозицій в певній мірі залежить від обраного способу угруповання. Отже, існує такий спосіб угруповання, при якому комерційні пропозиції Підприємства опиняються в кращому становищі щодо Конкурентів, ніж при інших способах.

Пропонуючи певний спосіб угруповання, підприємство впливає на оцінку Споживачем КП Конкурентів, тобто здійснює рефлексивне управління за схемою $S \rightarrow (C \rightarrow B)$ (рис. 4.14).

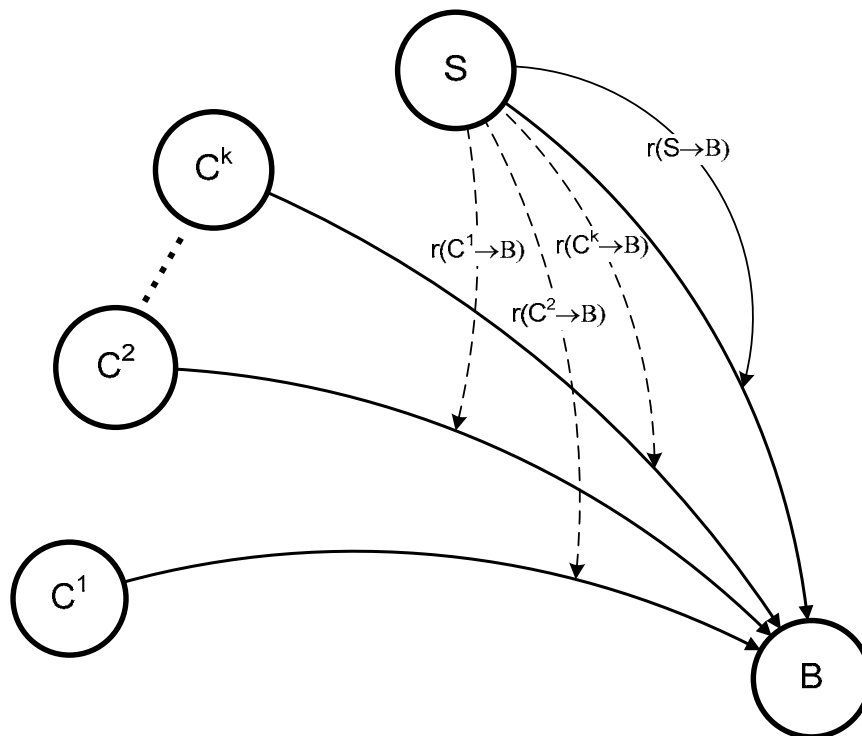


Рисунок 4.14 – Шляхи рефлексивного управління в системі Підприємство-Споживач-Конкуренти

Штрихові лінії, якими на рис. 4.14 показано вплив Керуючого суб'єкта на конкурентні пропозиції конкурентів відображають те, що ці впливи є непрямыми.

Розглянемо метод пошуку оптимального для Керуючого суб'єкта угруповання характеристик.

Нехай кожне з K конкуруючих КП можна описати за допомогою деякого набору з a_k характеристик. Нехай також можливо визначити L груп, на які можна розділити ці характеристики. Причому набір груп повинен відповідати наступним вимогам (В.4.1 – В.4.4):

В.4.1. Для L та всіх a_k виконується нерівність:

$$J < L < \min(a_k), \quad (4.5)$$

де J – мінімальна кількість груп, для якої, відповідно до [129], має виконуватися:

$$5 \leq J \leq 10. \quad (4.6)$$

В.4.2. Має існувати $n > 1$ способів розподілу всього набору характеристик по групах, таких, що задіяна кількість груп $g < L$, $g \rightarrow J$. При цьому кожна з характеристик повинна бути віднесена до однієї і тільки однієї групи.

В.4.3. З набору L може бути отриманий такий набір груп G , який би підходив для розподілу набору характеристик будь-якої КП. Причому має існувати $m > 1$ способів отримання набору G , з набору L .

В.4.4. Існує модель прийняття рішення про вибір той, чи іншої КП Суб'єктом управління. Очевидно, що дана модель не є зворотною, що характерно для функцій багатьох змінних.

Метою пошуку оптимальної угруповання є отримання такого набору G , при якому оцінка КП Керуючого суб'єкта відповідно до обраної моделі прийняття рішень, займає найвищу з можливих позицій в просторі вибору Суб'єкта управління.

Відповідно до В.4.1 – В.4.4, дана задача є NP-повною. Так, якщо в розглянутій задачі кожна з a_k характеристик може бути віднесена лише до двох різних груп, то за умови перебору всіх комбінацій необхідно оцінити 2^{a_k} варіантів розподілу [129, с. 98-100]. Для вирішення цієї задачі скористаємося інструментарієм генетичних алгоритмів.

Введемо наступні позначення:

ng – кількість груп, задіяних в хромосомі;

nk – кількість конкуруючих КП;

nl – кількість груп, які можуть бути використані для укрупнення множини характеристик КП;

nj – кількість найменувань характеристик, що описують всі конкуруючі КП;

$maxl$ – верхня межа оптимальної кількості груп;

$minl$ – нижня межа оптимальної кількості груп;

NA – множина найменувань характеристик КП;

NH – множина найменувань груп;

M_k – множина значень характеристик КП. Для визначеності можемо вважати, що множина M_0 відповідає Підприємству, а решта – конкурентам;

MR – множина, що задає ранги характеристик, тобто їх значення, щодо аналогічних характеристик інших КП;

A – матриця, що задає можливі відносини характеристика/група;

H – матриця відносин характеристика/група, що заповнюється на підставі хромосоми;

Z – множина, що задає корисність групових альтернатив;

R – матриця, яка містить рейтинги груп;

P – початкова хромосома, яка описує деякий варіант угруповання характеристик. Частини хромосоми, які відповідають окремим номерам груп мають назву «локуси», а можливі варіанти значень локусів насівають «алелі».

Для оптимальної роботи алгоритму бажано забезпечити потрапляння характеристик тільки в допустимі для них групи. Тому для виконання даної вимоги відносини (характеристика/група) задаються через матрицю A (табл. 4.13.), елементи якої a_{lj} приймають значення 1, якщо j -а характеристика може бути віднесена до l -й групи і значення 0 в іншому випадку.

Таблиця 4.13 – Структура матриці A

№ групи	№ характеристики					
	1	2	3	4	...	n_j
1	a_{11}	a_{12}	a_{13}	a_{14}		$a_{1 n_j}$
2	a_{21}	a_{22}	a_{23}	a_{24}		$a_{2 n_j}$
3	a_{31}	a_{32}	a_{33}	a_{34}		$a_{3 n_j}$
...						
nl	$a_{nl 1}$	$a_{nl 2}$	$a_{nl 3}$	$a_{nl 4}$		$a_{nl n_j}$

Перетворення алелей локусів хромосоми в номер групи здійснюється на стадії обчислення пристосованості хромосоми. Ця процедура здійснюється в кілька етапів:

1. Перейдемо в локусах хромосоми від інтервалу $(0; 1)$ до інтервалів $[1; j_{max}]$, де j_{max} – кількість груп, до яких може бути віднесена j -а характеристик.

$$x_j = \uparrow \left(p_j \cdot \sum_{l=1}^{nl} a_{lj} \right), \quad j \in (1; n_j), \quad (4.7)$$

де знаком « $\uparrow$ » представлено оператор округлення вгору.

2. На підставі X побудуємо матрицю H . За структурою матриця H аналогічна матриці A , але на її вміст накладається така умова:

$$\sum_{l=1}^{nl} h_{l,j} = 1, \quad j \in (1, n_j). \quad (4.8)$$

Ця умова буде виконуватися, якщо матрицю H заповнювати відповідно до:

$$h_{l,j} = \begin{cases} 1; & \sum_{l=1}^l a_{lj} = x_j, \\ 0 & \end{cases} \quad (4.9)$$

$$j \in (1, n_j); l \in (1, nl).$$

3. На підставі матриці H можна визначити кількість груп, задіяних генетичним алгоритмом для розміщення всієї множини характеристик:

$$ng = \sum_{l=1}^{nl} \overline{ng}_l,$$

$$\overline{ng}_l = \begin{cases} 1; & \sum_{j=1}^{nj} h_{l,j} > 0 \\ 0. & \end{cases} \quad (4.10)$$

4. На підставі матриці H і множини MR , що містить рейтинги характеристик кожного КП щодо інших, визначимо рейтингові оцінки одержаних груп за кожним КП і заповнимо матрицю R :

$$r_{k,l} = \frac{\sum_{j=1}^{nj} h_{l,j} \cdot mr_{kj}}{\sum_{j=1}^{nj} h_{l,j}}, \quad (4.11)$$

$$k \in (1, nk); l \in (1, nl).$$

5. Визначимо пристосованість хромосоми. Для цього покажемо спочатку, як визначається пристосованість в окремому випадку двох конкуруючих КП.

Відповідно до моделі Рестла [134, с.179], можна знайти ймовірність вибору Споживачем комерційної пропозиції $k1$ з двох варіантів $k1$ и $k2$ – $w(k1, k2)$:

$$w(k1, k2) = \frac{\left(\sum_{l=1}^{nl} \begin{cases} z_l & ; (r_{k1,l} > r_{k2,l}) \\ 0 & ; (r_{k1,l} < r_{k2,l}) \end{cases} \right)}{\left(\sum_{l=1}^{nl} \begin{cases} z_l & ; (r_{k1,l} > r_{k2,l}) \\ 0 & ; (r_{k1,l} < r_{k2,l}) \end{cases} \right) + \left(\sum_{l=1}^{nl} \begin{cases} 0 & ; (r_{k1,l} > r_{k2,l}) \\ z_l & ; (r_{k1,l} < r_{k2,l}) \end{cases} \right)}, \quad (4.12)$$

$$k \in (1, nk); l \in (1, nl).$$

Для того щоб забезпечити виконання вимоги В.4.2, слід ввести функцію штрафу $f(ng)$, яка приймає нульове значення, якщо кількість задіяних груп знаходиться за межами інтервалу $(\min l, \max l)$:

$$f(ng) = \begin{cases} 0; & \min l \leq ng \leq \max l \\ \left| \frac{ng - \bar{l}}{nl - \bar{l}} \right|; & (\min l > ng) \cup (ng > \max l). \end{cases} \quad (4.13)$$

$$\bar{l} = \frac{\min l + \max l}{2}.$$

Таким чином, функцію пристосованості для двох КП можна записати у такий спосіб:

$$v = w(k1, k2) - f(ng). \quad (4.14)$$

Якщо кількість КП $k > 2$, виконується попарне зіставлення конкурентних пропозицій, тому обсяги обчислень зростають прямо пропорційно до їх кількості.

При підготовці даних ОПР формує множини NA , M_k , MR , NH , Z , матрицю A і коефіцієнти $maxl$, $minl$, nk , nl , nj . При цьому, можна виділити наступні особливості:

Множини NA , M_k заповнюються на підставі аналізу конкуруючих КП. В першу чергу проводиться вибірка найменувань характеристик (NA) з доступних ОПР джерел. Потім кожній характеристиці кожного КП зіставляється деяке значення. Для кількісних характеристик це може бути їх безпосереднє значення, а для якісних – експертна оцінка. Необхідно вибрати єдиний інтервал для оцінки якісних показників, наприклад $(0,1)$.

Для поліпшення роботи алгоритму, множину MR , що задає ранги характеристик, тобто їх значення, щодо аналогічних характеристик інших КП, доцільно сформувати на стадії підготовки вхідних даних. Існує два способи формування цієї множини – із пропорційним, та рівномірним заповненням [166].

Множина найменувань груп (NH) формується, виходячи з аналізу інформації про конкуруючі КП, традиції представлення даних в галузі, а також власної інтуїції ОПР. У формованих групах має бути присутня інформація про нормативну, науково-технічну, організаційну, комерційну та економічну категорії розглянутих КП [166].

Визначення значущості кожної з груп, тобто компонентів множини Z , проводиться експертним шляхом.

Матриця A заповнюється на основі логічного зіставлення даних, що містяться у множинах NA і NH . Якщо j -я характеристика може входити в l -ї групи, то елементу a_{lj} присвоюється значення 1. В іншому випадку елементу a_{lj} присвоюється значення 0.

Нарешті, константи nk , nl , nj задаються відповідно до розмірності множин M_k , NA , NH .

Таким чином, вирази (4.7) – (4.14) є моделлю визначення пристосованості деякого припустимого рішення задачі, закодованого у вигляді хромосоми генетичного алгоритму. Розглянемо можливості реалізації цієї моделі в пакеті *Evolver* (рис.4.15).

Генетична модель (рис. 4.15) задається стандартними засобами *Microsoft Excel*. Використано умовні вхідні дані. Незважаючи на те, що в *Evolver* не накладено істотних обмежень на вид і діапазон оптимізуються значень, досвід роботи з цим пакетом дозволяє зробити висновок про те, що оптимізація проводиться швидше і ефективніше, якщо комірок, які оптимізуються, небагато

і вони максимально наближені до представлення у вигляді хромосоми [147]. Оскільки дана модель спочатку будувалася в розрахунку на генетичну оптимізацію, дану умову в неї виконано, і відповідність кожної характеристики до набору груп задається одним числом, що знаходиться в інтервалі (0; 1).

	A	B	C	D	E	F	G	H	I	J	K	L	M
1													
2		Матриця А											
3		1	2	3	4	5	6	7	8	9	Сумма	Хромосома	Х
4	1	1	1				1	1	1		5	0,529291446	3
5	2		1			1			1		3	0,159077565	1
6	3			1	1			1			3	0,633964727	2
7	4	1	1				1		1	1	5	0,941491909	5
8	5			1				1			2	0,718110933	2
9	6	1		1		1		1	1	1	6	0,600417651	4
10	7		1		1			1			3	0,680189485	3
11	8						1		1		2	0,441993449	1
12	9					1	1			1	3	0,000354218	1
13	10		1		1			1			3	0,329114486	1
14	11	1		1			1		1		4	0,735652778	3
15													
30		Матриця Н										Рейтинг характеристики	
31		1	2	3	4	5	6	7	8	9		Пр-тие	Конк 1
32	1	0	0	0	0	0	1	0	0	0		0,44	0,56
33	2	0	1	0	0	0	0	0	0	0		0,49	0,51
34	3	0	0	0	1	0	0	0	0	0		0,92	0,08
35	4	0	0	0	0	0	0	0	0	1		0,72	0,28
36	5	0	0	0	0	0	0	1	0	0		0,07	0,93
37	6	0	0	0	0	0	0	1	0	0		0,43	0,57
38	7	0	0	0	0	0	0	1	0	0		0,07	0,93
39	8	0	0	0	0	0	1	0	0	0		0,97	0,03
40	9	0	0	0	0	1	0	0	0	0		0,99	0,01
41	10	0	1	0	0	0	0	0	0	0		0,52	0,48
42	11	0	0	0	0	0	1	0	0	0		0,50	0,50
43	Задействовано	0	1	0	1	1	1	1	0	1	6		
44	Сумма	0	2	0	1	1	3	3	0	1			
45	R0	0,00	0,51	0,00	0,92	0,99	0,64	0,19	0,00	0,72			
46	R1	0,00	0,49	0,00	0,08	0,01	0,36	0,81	0,00	0,28			
47													
48	Полезность групповых альтернатив	20	25	15	18	31	11	19	5	27			
49													
50	Пр-тие	0	25	0	18	31	11	0	0	27	112		
51	Конкурент	0	0	0	0	0	0	19	0	0	19		
52													
53	Вероятность выбора КП пр-тия	0,854962											
54													
55	Штраф	0		1		min_I	5	max_I	10	груп вс	9		7,5
56													
57	Приспособленность	0,855											
58													

Рисунок 4.15 – Реалізація моделі оптимізації способу простору вибору

Після опису базової моделі засобами *Excel*, необхідно задати набір комірок, що змінюються, набір обмежень на проміжні значення і цільову комірку, яка містить значення функції пристосованості (рис. 4.16).

Після установки параметрів моделі можна перейти до установки параметрів генетичного алгоритму (рис. 4. 17).

На вкладках (рис. 4.17) відкривається доступ до таких груп параметрів:

параметри оптимізації – розмір популяції, спосіб генерації стартової популяції, ймовірність скоєння кросоверу та мутації і тому подібні;

умови припинення роботи алгоритму – кількість поколінь, час рахунку, відсутність прогресу, виконання математично заданої умови, помилка рахунку, та їх комбінації;

параметри візуалізації роботи алгоритму – дозволяють, чи забороняють динамічне відображення зміни функції пристосованості алгоритму;

виконання заданих користувачем макрокоманд на різних етапах роботи алгоритму.

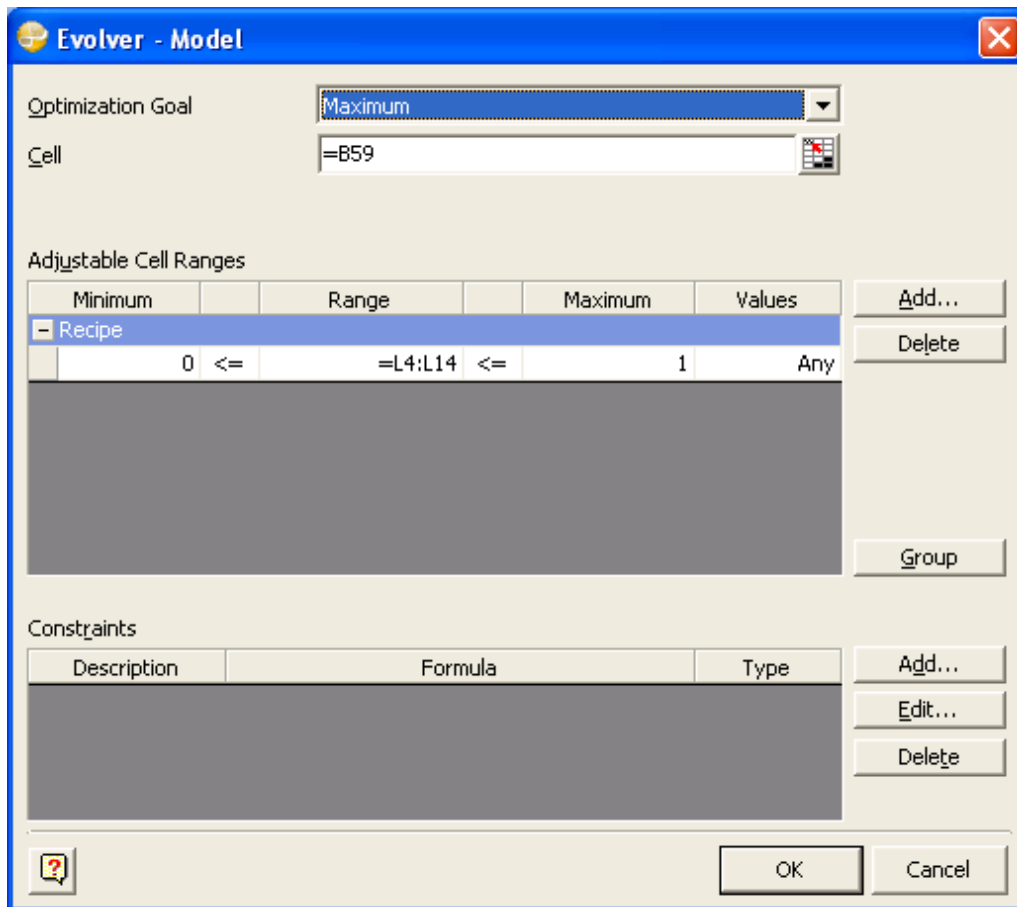


Рисунок 4.16 – Установка параметров генетической модели

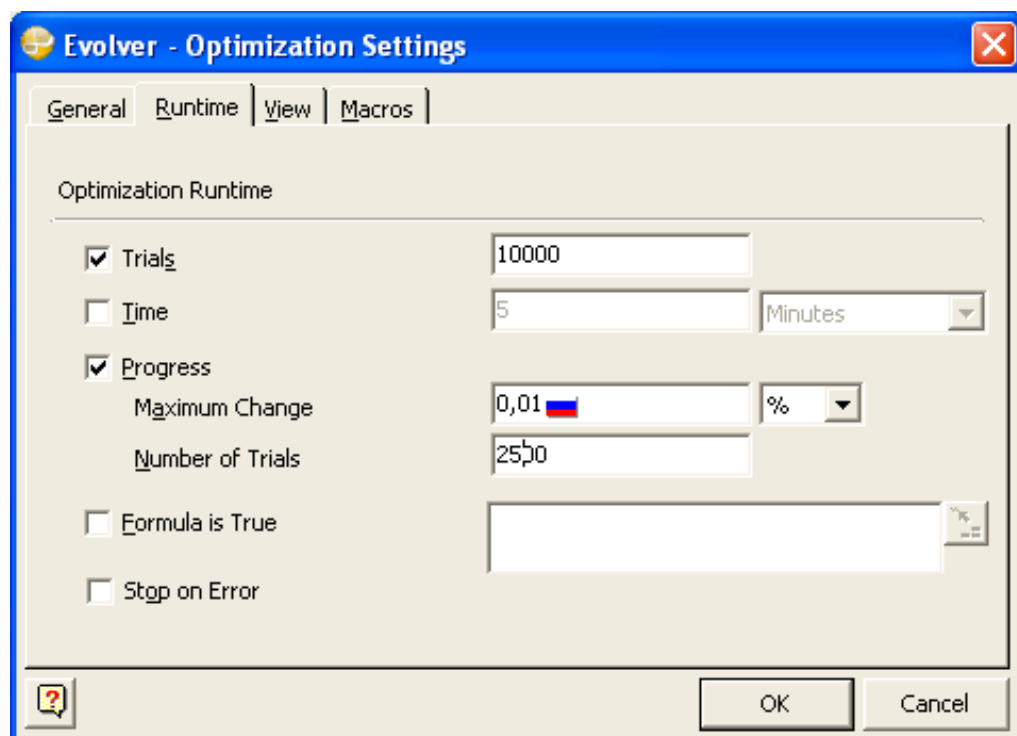


Рисунок 4.17 – Установка параметров генетического алгоритма

Оскільки точна настройка параметрів генетичного алгоритму необхідна тільки для роботи зі складними просторами рішень, в даному випадку значення більшої частини налаштувань залишено без змін. В якості умов зупинки алгоритму запропоновано такі: досягнення 10000 поколінь, або відсутність істотних збільшень функції пристосованості протягом 2500 поколінь.

Результати моделювання

При оптимізації моделі (рис 4.15) в *Evolver* роботу генетичного алгоритму було припинено після 3265 ітерацій. Причому максимальне значення було досягнуто вже на 764-й ітерації. Зупинка роботи алгоритму настала після виконання умови на відсутність істотного поліпшення значень цільової функції протягом 2500 ітерацій. Час розрахунків на комп'ютері з процесором Intel Atom N450 (1 ГГц) склав 31 секунду.

На рис. 4.18 показано графік зміни значень цільової функції для кращої хромосоми в популяції.

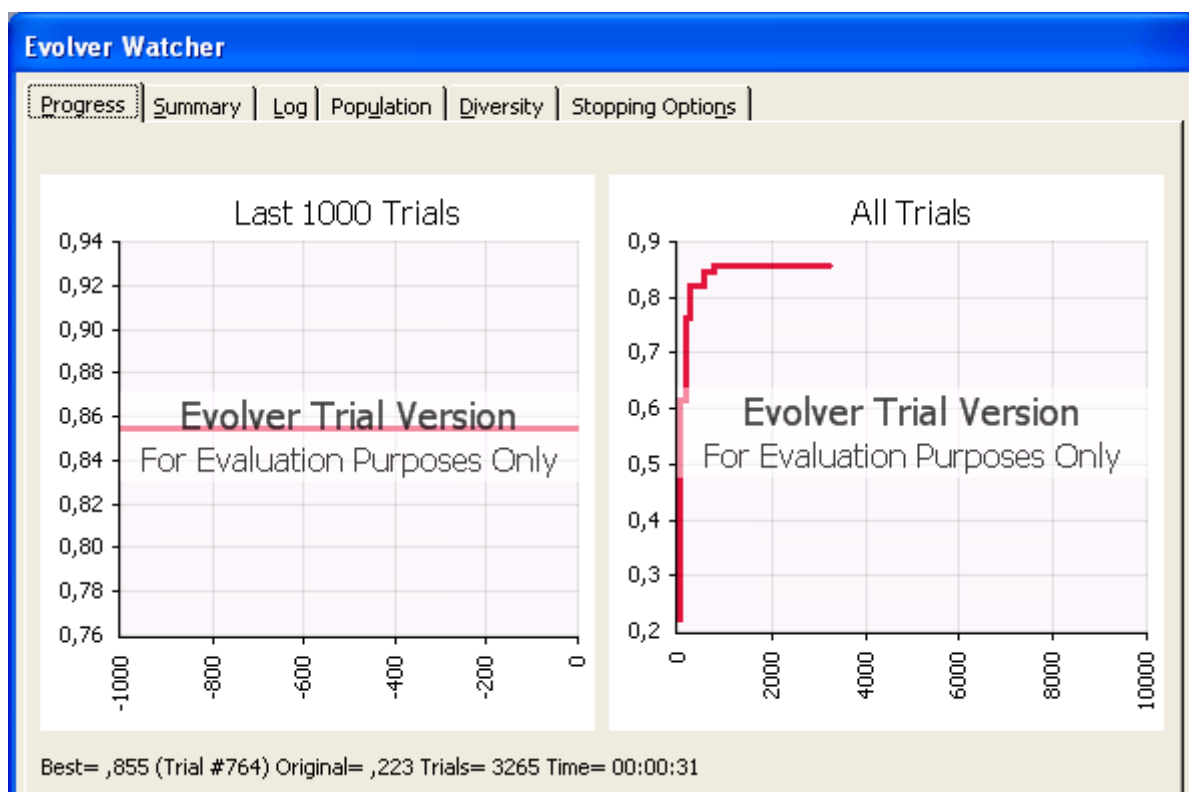


Рисунок 4.18 – Зміна функції пристосованості в процесі роботи ГА

Інформація, що наведена на рис. 4.18 в графічному вигляді, служить лише для наочного відображення прогресу роботи ГА. Детальна інформація міститься в лог-файлах роботи алгоритму. Так, на рис. 4.19 показано вікно програми, що містить числові дані про прогрес значень цільової функції і відповідні їм значення хромосомного набору (вхідних змінних моделі).

З аналізу рис. 4.19 видно, що найвищі темпи збільшення значень цільової функції характерні для початкових етапів роботи алгоритму. Так вже до 277 ітерації значення ймовірності вибору КП підприємства з початкових 0,223 зросла до 0,82, що всього на 4% менше максимально досягнутого після

закінчення роботи алгоритму значення – 0,855. В цілому, аналіз динаміки зміни значень функції пристосованості (рис. 4.18, рис. 4.19) дозволяє зробити висновок про те, що простір рішень задачі є плавним і добре підходить для роботи генетичного алгоритму.

Trial	Elapsed Time	Result	L4	L5	L6	L7	L8
1	00:00:01	,223	,3813	,3683	,4875	,0283	,2961
2	00:00:02	,287	,3813	,218	,4875	,0283	,2556
7	00:00:02	,3525	,3813	,3683	,4875	,0283	,2961
12	00:00:03	,3597	,3813	,3683	,6916	,0283	,2961
14	00:00:03	,446	,3813	,3683	,4875	,0283	,2961
27	00:00:03	,6148	,5293	,0004	,4875	,0283	,2961
79	00:00:03	,6694	,3813	,3683	,8438	,0283	,2961
111	00:00:04	,7556	,5293	,0004	,8934	,0283	,2961
123	00:00:04	,7626	,1809	,3683	,8438	,0283	,2961
272	00:00:05	,774	,5293	,3683	,8438	,895	,2961
277	00:00:05	,8208	,5293	,0004	,8438	,0283	,8877
494	00:00:07	,843	,5293	,1591	,2752	,0283	,8877
512	00:00:07	,8468	,6938	,3683	,639	,0283	,8877
764	00:00:09	,855	,5293	,1591	,634	,9415	,7181

Рисунок 4.19 – Збільшення значень функції пристосованості (за поколіннями)

Результати оптимізації генетичної моделі показані в табл. 4.14.

Таблиця 4.14 – Результати генетичної оптимізації способу простору вибору

Характеристика	Група								
	1	2	3	4	5	6	7	8	9
Характеристика 1	0	0	0	0	0	1	0	0	0
Характеристика 2	0	1	0	0	0	0	0	0	0
Характеристика 3	0	0	0	1	0	0	0	0	0
Характеристика 4	0	0	0	0	0	0	0	0	1
Характеристика 5	0	0	0	0	0	0	1	0	0
Характеристика 6	0	0	0	0	0	0	1	0	0
Характеристика 7	0	0	0	0	0	0	1	0	0
Характеристика 8	0	0	0	0	0	1	0	0	0
Характеристика 9	0	0	0	0	1	0	0	0	0
Характеристика 10	0	1	0	0	0	0	0	0	0
Характеристика 11	0	0	0	0	0	1	0	0	0
Характеристик в групі	0	2	0	1	1	3	3	0	1
Рейтинг підприємства	0,00	0,51	0,00	0,92	0,99	0,64	0,19	0,00	0,72
Рейтинг конкурента	0,00	0,49	0,00	0,08	0,01	0,36	0,81	0,00	0,28
Корисність групи	20	25	15	18	31	11	19	5	27

З аналізу табл. 4.14 видно, що за результатами проведеної оптимізації алгоритм розташував вихідні характеристики в 6 груп з 9 можливих. При цьому тільки по групі №7 продукція конкурента виявляється краще, ніж продукція підприємства.

Таким чином, задача формування та оптимізації рефлексивних впливів, при виконанні певних умов на область припустимих рішень, може бути формалізована за використанням сучасного інструментарію генетичного моделювання, що дозволяє істотно знизити вимоги до кваліфікації фахівців і розширити застосування рефлексивного управління в економіці країни.

Отже розглянуті аспекти реалізації інноваційних моделей і методів для прийняття рішень в економіці дозволили обґрунтувати вибір програмного забезпечення для нейромережевого та генетичного моделювання. Із використанням методологічних положень, викладених в р.1 – р.3 розроблено ефективні рішення задач прогнозування показників бюджетного процесу, а також формування та оптимізації рефлексивних впливів промислових підприємств.

ЗАГАЛЬНІ ВИСНОВКИ

Модернізація систем управління економічними системами має за мету підвищення ефективності суспільного виробництва і подолання цифрових розривів. Тільки інноваційний розвиток внутрішніх економічних процесів може вивести вітчизняну економіку з кризи. Тому впровадження сучасних інтелектуальних методів аналізу, обробки інформації і прийняття рішень є одним з першочергових завдань, які ставляться економікою країни перед українськими підприємствами та організаціями.

Розглянуті в монографії методологічні підходи, моделі й методи спрямовані на вдосконалення процесів підготовки та прийняття рішень в економічних системах всіх рівнів за рахунок використання інтелектуальних методів для аналізу й обробки даних, розробки рішень та аналізу їх ефективності. Запропонована методологія моделювання інноваційних інтелектуальних систем прийняття рішень створює передумови підвищення адаптивності та життєздатності економічних систем.

Проведений аналіз дозволив виділити тенденції розвитку інформаційного середовища, характерні для сучасного суспільства. З розвитком і поширенням обчислювальної техніки, відбувається різке збільшення її ролі в процесах прийняття рішень від суто обчислювальних задач і забезпечення швидкого доступу до даних для осіб, що приймають рішення до повноважень в самостійному прийнятті рішень в умовах невизначеності. Використання інтелектуальних методів аналізу слабоструктурованих даних для прийняття рішень в складних системах дозволяє знизити витрати на розробку і супровід систем управління, підвищити їх точність, та є актуальним для вирішення більшості економічних задач, які можуть бути описані за формальними ознаками і не засновані на використанні суто людських якостей.

Бурхливий розвиток інформаційного суспільства, який зараз спостерігається, спричинив виникнення нових задач для інтелектуального аналізу та обробки даних. Швидкість цих процесів є такою, що існуючі підходи до класифікації відповідних економічних задач вже не відповідають фактичному стану. У той же час систематизація та класифікація є одним з найважливіших компонентів наукових досліджень, який сприяє розвитку науки і переходу її з емпіричного на системний рівень. Проведені дослідження дозволили вдосконалити та актуалізувати класифікацію економічних задач з аналізу й обробки даних та інтелектуальних методів їх вирішення.

Формалізація процесів синтезу є важливим компонентом методології моделювання складних інформаційних систем. Проведений аналіз логічної і часової структури, яка характерна для інноваційних інтелектуальних систем прийняття рішень, дозволив визначити необхідність використання в цьому випадку структурно-параметричного підходу. Структурний синтез дозволяє визначити оптимальний набір методів вирішення поставлених завдань, і їх послідовність. Параметричний дозволяє оптимізувати елементи системи. Запропонований методологічний підхід який засновано на морфологічному методі Ф. Цвіккі та використанні апарату n-дольних гіперграфів дозволяє

формалізувати процеси синтезу інноваційних інтелектуальних систем прийняття рішень та урахувати обмежену придатність методів для вирішення різних класів економічних задач. Використання цього підходу дозволяє знизити витрати на процес розробки інтелектуальній системи прийняття рішень

Завершальною частиною індуктивної фази дослідження є сформована авторська концепція, що відображає суть дослідження. Концепція визначає стратегію дій в рамках наступної дедуктивної фази дослідження, тобто його розвитку від абстрактного до конкретного. Саме на підставі концепції виділяємо комбінацію моделей і методів, що забезпечують досягнення мети дослідження. В запропонованій концепції процес прийняття рішень зводиться до сукупності процесів спостереження, моделювання, ідентифікації, оцінювання та вибору. Реалізація цих процесів базується на розробках вітчизняних і зарубіжних авторів в таких областях, як статистичний аналіз, спостереження і експеримент, технології баз даних, математичне моделювання, інтелектуальних аналіз і обробка даних, методи оптимізації. Концепція моделювання інноваційної інтелектуальної системи прийняття рішень дозволяє звести задачу пошуку і прийняття рішень до набору більш простих завдань, вирішення яких може бути здійснено з використанням інтелектуальних обчислень, що дозволяє підвищити ефективність прийнятих рішень, скоротити витрати на підтримку роботи системи в умовах мінливого зовнішнього середовища, зменшити вимоги до персоналу, і витрати на його навчання, підвищити автономність роботи і життєздатність системи.

Розвиток методології моделювання штучних нейронних мереж для вирішення економічних задач в монографії проведено за напрямками вибору інструментів моделювання, визначення архітектури штучних нейронних мереж, вдосконалення методів роботи із вхідними даними, дослідження ефективності.

Доведено, що ефективність застосування інтелектуальних методів прийняття рішень безпосередньо пов'язана з постановкою завдань. Визначено поняття базової постановки завдання. На прикладі вирішення задачі створення біржової торговельної системи та системи прогнозування фінансового стану комерційних банків України визначено необхідність зведення розв'язуваної задачі до однієї або декількох базових постановок та визначення економічної ефективності реалізації кожної з них. Остаточний вибір методів та інструментів вирішення конкретної задачі відбувається на підставі аналізу результатів, отриманих при її вирішенні в різних базових постановках.

Проблеми визначення оптимальної архітектури штучних нейронних мереж, визначення достатнього обсягу навчальної вибірки і відбору вхідних даних знаходяться в тісному взаємозв'язку. Актуальність цих проблем впливає з відомої залежності між складністю структури нейронної мережі і кількістю прикладів, які необхідні для її навчання, що особливо проявляється при нестачі вхідних даних. Комплексне дослідження даної проблеми дозволило систематизувати теоретичні висновки щодо формалізації відношень між такими параметрами, як розмірність вектору вхідних даних, кількість ступенів свободи нейронної мережі, та обсяг вхідної вибірки. Дефіцит обсягу вхідної вибірки, або її різноманітності є стримуючим фактором для досягнення високих

показників ефективності нейромережевого моделювання. Запропоновано методи управління розмірністю даних. Для її зниження розглянуто використання логічного аналізу даних, методів оптимізації їх представлення, прямих та непрямих методів аналізу значущості. При малій розмірності вектора вхідних даних, великій кількості прикладів у вибірці і наявності складної залежності між вхідними і вихідними параметрами актуальним є завдання підвищення різноманітності, для вирішення якого запропоновано використання результатів емпіричних методів аналізу та перехід до концепції нейро-експертної архітектури.

При розгляді методів вирішення економічних задач природним чином виникає питання дослідження їх ефективності. Особливості технологій машинного навчання зумовлюють високу варіативність отриманих результатів, що не дає можливість заздалегідь оцінити вигоди, які дає їх застосування. Проблему аналізу ефективності інтелектуальних методів вирішення економічних задач розглянуто по відношенню до різних об'єктів і процесів, серед яких алгоритми і програмне забезпечення, процес навчання, а також результати аналізу, або обробки даних.

Для зіставлення алгоритмів машинного навчання і їх реалізації, запропоновано концепцію типових задач, основні характеристики яких є досить близькими для деякого набору інших задач у схожій постановці. Висунуто вимоги до задач, які можна використовувати в якості типових та розглянуто їх приклади. Зіставлення ефективності програм і алгоритмів інтелектуальних обчислень при вирішенні типових задач дає змогу їх апіорного зіставлення, що знижує витрати на реалізацію інноваційних інтелектуальних систем прийняття рішень та підвищує їх ефективність.

Систематизовано та розширено підходи до оцінки результатів інтелектуальних обчислень. Виділено три типи моделей аналізу даних за ознакою метрик оцінки економічної ефективності – кількісні моделі першого і другого типів та дескриптивні моделі. Систематизовано методи оцінки ефективності для кожного з цих типів. Досліджено існуючі підходи до аналізу ефективності методів обробки даних та запропоновано метод оцінки, який дозволяє забезпечити пріоритетну значимість правильного ранжирування об'єктів з початку списку, що краще відповідає умовам економічних задач з ранжирування.

Для здійснення обґрунтованого вибору інструментальних засобів реалізації інтелектуальних обчислень порівняння має здійснюватися за багатьма критеріями, які до того ж мають різну значущість. Для вирішення даного завдання розроблено нечітку модель оцінки та формалізовано процедури фазифікації вхідних даних і дефазифікації результатів. Підставою для формування множини критеріїв оцінки є вимоги стандарту ISO 9126-4:2004 «Характеристики та метрики якості програмного забезпечення». Нечітко-логічна модель може ефективно працювати при великій кількості аналізованих продуктів і критеріїв, а також дозволяє врахувати всі параметри досліджуваних продуктів. Для випадків, коли оцінки декількох програмних продуктів виявилися близькими одна до одної, в роботі запропоновано графічний метод

зіставлення, заснований на використанні павутинних діаграм, які наочно показують співвідношення переваг і недоліків окремих інструментальних засобів. Отже павутинні діаграми можуть використовуватися для остаточного прийняття рішень з вибору, після попереднього відбору за допомогою нечіткої логічної моделі.

Проведено вирішення задачі біржового спекулянта у регресійній, класифікаційній та кластеризаційній постановках. Проведене дослідження показує, що технології моделювання для регресійній та класифікаційній постановок мало відрізняються одна від одної. Основна різниця виявилася в інтерпретації отриманих даних. Однак, класифікаційна модель, забезпечує найкращі питомі показники, ніж регресійна, отже вона більше підходить для створення автоматичних торговельних систем. У той же час регресійна модель, як і кластеризаційна, забезпечує отримання прогнозу в кожному біржовому періоді, тому добре підходить для створення автоматизованих систем підтримки прийняття рішень. Відносно низькі питомі показники прибутковості було отримано в результаті рішення задачі біржового спекулянта у регресійній і кластеризаційній постановках. Значною мірою це обумовлено застосуванням спрощеного методу прийняття торгових рішень, який не враховує силу прогнозованих рухів валютного курсу. Однак якщо ці моделі застосовуються для створення систем підтримки прийняття рішень, то цей недолік не має принципового значення, оскільки мозок людини – трейдера здатний фільтрувати інформацію про малоймовірні зміни курсу.

При вирішенні задачі прогнозування банкрутств комерційних банків більш вдалою виявилася її постановка, як задачі кластеризації, тобто угруповання об'єктів. В рамках цієї постановки використання самоорганізаційних нейронних мереж дозволяє отримати достатньо достовірні оцінки, які можна використовувати в практичній діяльності для визначення надійності потенційних фінансових партнерів. Хоча отримання абсолютно достовірного прогнозу в поточних економічних і політичних умовах не є можливим, запропонований метод може відігравати роль одного з індикаторів фінансової стійкості при виборі банків, страхових компаній та інших фінансових посередників. В цілому, отримані результати підтвердили гіпотезу про те, що результат рішення задачі інтелектуальних обчислень безпосередньо пов'язаний з її постановкою.

Темпи накопичення інформації в сучасному суспільстві перевищують темпи розвитку засобів та методів її обробки. Як один з компонентів даної проблеми вирішується проблема спрощення даних, представлених у вигляді динамічних рядів. На підставі аналізу існуючих підходів спрощення даних, запропоновано концепцію квантування за часом зі змінним кроком, для реалізації якої розроблено математичну модель та реалізовано її із використанням інструментарію генетичних алгоритмів. Запропонований метод дозволяє швидко і ефективно скорочувати кількість точок відліку ряду з будь-яким ступенем стиснення даних, який регулюється зміною лише одного показника. У порівнянні з існуючими методами квантування він забезпечує збереження пікових значень ряду, а також більш ефективне стиснення даних,

особливо якщо значення показників ряду змінюються повільно. Квантування за часом зі змінним кроком може використовуватися в економічній статистиці і аналізі, в процесі підготовки презентаційних матеріалів, а також у всіх інших випадках, коли потрібно виділити основні тенденції у великій послідовності даних.

Непрямі методи оцінки параметрів економічних систем дозволяють розширити різноманітність вхідних даних для аналізу та підвищити обґрунтованість прийнятих рішень. Розроблено методи оцінки параметрів зовнішніх та внутрішніх факторів кредитних ризиків комерційних банків із використанням системно-динамічних імітаційних моделей. Модель ціноутворення на ринку житлової нерухомості ґрунтується на моделях причинно-наслідкового взаємозв'язку між попитом, пропозицією і ціною, а також між факторами, що формують фінансові можливості покупців на ринку нерухомості. Вона може бути використана для аналізу взаємозв'язків між факторами ринку і для прогнозування. В роботі проведено аналіз впливу зовнішніх факторів на ціни ринку житлової нерухомості в умовах України. Імітаційні експерименти показали адекватну відповідність результатів моделювання реальним змінам цін на нерухомість, а також дозволили проаналізувати альтернативні сценарії розвитку цього ринку.

Моделювання внутрішніх факторів ризику банківських активних операцій здійснено на прикладі процесів реструктуризації кредитів. Одним з найтяжчих наслідків кризових явищ в економіці для українських банків слід вважати надмірне зростання простроченої кредитної заборгованості, обсяг якої в десятки разів перевищує норми, які прийняті в світовій практиці, та вимірюється вже сотнями мільярдів гривень. Боротьба з ризиками виникнення неплатежів по кредитним угодам, таким чином безумовно є актуальним завданням для банківської системи України. Запропонована модель дозволяє проаналізувати зміну грошових потоків сімейного бюджету при зміні балансу прибутків і витрат, а також розрахувати можливі сценарії реструктуризації простроченої кредитної заборгованості. Впровадження запропонованої схеми реструктуризації дозволяє уникнути ризику виникнення простроченої заборгованості за кредитним договором.

Порівняльний аналіз інструментальних засобів моделювання нейронних мереж серед програмного забезпечення низькорівневої розробки, програмованих систем математичного моделювання та інтерактивних програмних платформ дозволив виявити програмні продукти, які за комплексом характеристик є зручнішими на етапах для попереднього аналізу даних і моделювання.

Аналіз програмного забезпечення для генетичного моделювання проводився з позицій вартості програмних продуктів, зручності у користуванні, та швидкодії, яку було протестовано на прикладі задачі комівояжера – типової для оптимізаційних алгоритмів. Виявилось, що кожна з розглянутих програм має власні переваги та недоліки. Низькорівневе програмне забезпечення відрізняється високою швидкістю, та безкоштовністю, але складне для використання. Спеціальне програмне забезпечення яке розглянуто на прикладі

пакета MatLab оптимально в навчальних та наукових організаціях, а також для розробки комерційних продуктів. Вбудоване програмне забезпечення є найбільш простим в освоєнні і використанні та дозволяє вирішувати задачі в режимі, близькому до інтерактивного, але має найнижчу швидкодію.

Обране програмне забезпечення застосовано для вирішення практичних завдань. Нейромережеве моделювання процесів виконання державного бюджету із використанням запропонованих в монографії методів роботи із вхідними даними, визначення архітектури нейронної мережі, та аналізу ефективності процесу навчання дозволило скласти прогноз, який в подальшому було підтверджено фактичними даними. Результати прогнозування показників виконання бюджетного процесу можуть бути використані виконавчими і законодавчими органами влади для прийняття рішень при формуванні і здійсненні фінансової політики, виявлення резервів для залучення коштів у бюджет, підвищення ефективності їх використання, посилення контролю за їх формуванням і витратами.

Для реалізації моделей генетичної оптимізації рефлексивних впливів промислових підприємств застосовано вбудоване програмне забезпечення. Рефлексивне управління є подальшим розвитком ігрового підходу до аналізу процесів прийняття рішень та дозволяє враховувати не тільки об'єктивні фактори, а й схильності суб'єктів гри та їх бачення ситуації. Будучи зручним та ефективним для таких сфер, як міждержавні взаємини, зовнішня і внутрішня політика, військова справа, використання даного інструментарію для аналізу та управління економічними взаєминами між підприємствами утруднено через низку обставин, серед яких слабка поляризація альтернатив у об'єкта рефлексивного управління. Запропонована модель дозволяє підсилити поляризацію, та забезпечити для комерційних пропозицій керуючого суб'єкта краще становище відносно конкурентів. Однак, така модель має ознаки NP-повноти, тому для її оптимізації використано генетичний алгоритм, який дозволив знайти оптимальне рішення.

Таким чином, розроблена методологія моделювання інноваційних інтелектуальних систем прийняття рішень в економіці вирішує нагальну для управління економікою України проблему подолання цифрових розривів та надає можливість вирішити проблеми підвищення ефективності функціонування економічних систем різного рівня, що відкриває нові можливості для економічного росту.

- 1 Allcott H. Social Media and Fake News in the 2016 Election / Hunt Allcott, Matthew Gentzkow // Journal of Economic Perspectives – Vol. 31, № 2, Spring 2017 – Pp 211–236.
- 2 Alter S. A. Decision Support Systems: Current Practice and Continuing Challenges / S. A Alter. – Reading, Mass.: Addison – Wesley Publ. Co., 1980. – 316 p.
- 3 Around the World in Dollars & Cents / What Price The World? [Electronic resource]. – Available at: \www/URL: http://www.savills.co.uk/research_articles/188297/198669-0/
- 4 Baum E. B. What size net gives valid generalization? / E. B. Baum D. Haussler.// Neural Computation, 1989, vol. 1, Pp. 151-160.
- 5 Breiman Leo. Classification and regression trees / Leo Breiman, J. H. Friedman, R. A. Olshen, C. J Stone. – Monterey, CA: Wadsworth & Brooks/Cole Advanced Books & Software. 1984. – 368 p.
- 6 Bremermann H. J. Optimization through evolution and recombination / Hans J. Bremermann // In: Yovits M. C., et al. (eds.) Self-Organizing Systems, Washington, Spartan Books, 1962. – Pp. 93–106.
- 7 Brooks C. Real Estate Modeling And Forecasting / Chris Brooks, Sotiris Tsolacos. – New York: Cambridge University Press, 2010. – 453 p.
- 8 Chapman P. “CRISP-DM 1.0: Step-by-step data mining guide / P. Chapman, J. Clinton, R. Kerber, T. Khabaza, T. Reinartz, C. Shearer, R. Wirth – NCR Systems Engineering Copenhagen (USA) and DaimlerChrysler AG (Germany), SPSS Inc. (USA) and OHRA Verzekeringen Bank Denmark), Group B.V (The Netherlands), 2000.
- 9 Cho D. Dynamic Relationship between Housing Values and Interest Rates in the Korean Housing Market / D. Cho, S. Ma // The Journal of Real Estate Finance and Economics, Vol. 32, 2006. – Pp.180-183.
- 10 Comparison of deep learning software [Electronic resource]. – Available at: \www/URL: https://en.wikipedia.org/wiki/Comparison_of_deep_learning_software
- 11 Comparison of Neural Network Simulators [Electronic resource]. – Available at: \www/URL: https://grey.colorado.edu/emergent/index.php/Comparison_of_Neural_Network_Simulators
- 12 Copeland J. Artificial Intelligence: A Philosophical Introduction / Jack Copeland – Wiley-Blackwell. September 1993. – 328 p.
- 13 Crouhy M. A Comparative analysis of current credit risk models / M. Crouhy, D. Galai, R. Mark // Journal of Banking and Finance, 2000. – Vol. 24.– No. 1-2. – Pp. 59-117
- 14 Davenport T. H. Competing on Analytics: The New Science of Winning / T. H. Davenport; J. G. Harris – Harvard: Business School Press, 2007. – 240 p.
- 15 De Jong K. A. Evolutionary Computation: a unified approach / Kenneth A. De Jong – MIT Press, 2006. – 179 p.

- 16 Eskinasi M. Towards housing system dynamics: Projects on embedding system dynamics in housing policy research. Dissertation thesis. / M. Eskinasi. – Eburon: Academic Publishers, 2014. – 165 p.
- 17 Evolution of data mining, Gartner Group Advanced Technologies and Applications Research Note, 2/1/95. [Electronic resource]. – Available at: \www/URL: <http://www.thearling.com/text/dmwhite/dmwhite.htm>
- 18 Fedorowicz J., Document Based Decision Support in Decision Support for Management / J. Fedorowicz , in R. Sprague Jr. and Hugh J Watson (eds.) – Upper Saddle River, N.J.: Prentice-Hall, 1996. – 371 p.
- 19 Forrester J. W. Urban Dynamics / J. W. Forrester – Cambridge MA: MIT Press, 1969. – 299 p.
- 20 Ginzberg M. J. Decision Support Systems: Issues and Perspectives / M. J. Ginzberg , E. A. Stohr // Processes and Tools for Decision Support. / Ed. by H. G. Sol. – Amsterdam: North-Holland Publ. Co., 1983. – Pp. 9-31.
- 21 Guo K. L. DECIDE: a decision-making model for more effective decision making by health care managers / Kristina L. Guo // The Health Care Manager. №27 (2), June 2008. – Pp. 118–127.
- 22 Haettenschwiler P. Neues anwenderfreundliches Konzept der Entscheidungsunterstützung. Gutes Entscheiden in Wirtschaft, Politik und Gesellschaft / P. Haettenschwiler // Zurich: Hochschulverlag AG, 1999. – S. 189–208.
- 23 Han J. Data Mining: Concepts and Techniques / Jiawei Han, Micheline Kamber. – 2nd Edition, Morgan Kaufmann, 2006. – 761 p.
- 24 Han J. Data Mining: Concepts and Techniques / Jiawei Han, Micheline Kamber, Jian Pei. – 3rd Edition, Morgan Kaufmann, 2012. – 740 p.
- 25 Hilbert M. The World's Technological Capacity to Store, Communicate, and Compute Information / Martin Hilbert , Priscila López // Science 01 Apr 2011: Vol. 332, Issue 6025. Pp. 60-65
- 26 Holland J. H. Adaptation in Natural and Artificial Systems: An Introductory Analysis with Applications to Biology, Control and Artificial Intelligence / John H. Holland – MIT Press, 1975. – 225 p.
- 27 Holland J. H. Genetic algorithms and the optimal allocation of trials / John H. Holland // SIAM Journal on Computation, 2: 1973, Pp. 88-105.
- 28 Hothorn T. Unbiased Recursive Partitioning: A Conditional Inference Framework / T. Hothorn, K. Hornik, A. Zeileis // Journal of Computational and Graphical Statistics, 15(3), 2006. – Pp. 651–674
- 29 Human Branch Project. Overview. [Electronic resource]. – Available at: \www/URL: <https://www.humanbrainproject.eu/>
- 30 Hunt E. B. Experiments in Induction / E. B. Hunt, J. Marin, P. T. Stone // New York, Academic Press, 1966, V. 1. pp. 45–69
- 31 Hwang S. Korean Real Estate Market Mechanisms and Deregulation of Mortgage Loans: Qualitative Analysis / S. Hwang, M. Park, H. S. Lee // 27th International System Dynamics Conference, Albuquerque NM, 2009. – 11p.
- 32 Karpinski M. Polynomial bounds for VC dimension of sigmoidal and general Pfaffian neuronal networks / M. Karpinski , A. Macintyre. // Journal of Computer and System Sciences, vol.54, 1997. – Pp. 169-176.

- 33 Kohonen T. Self-Organizing Maps (Third Extended Edition) / T Teuvo Kohonen. – New York, 2001. – 501 p.
- 34 Koiran P. Neural networks with quadratic VC dimension / P. Koiran, E. D. Sontag. // Advances in Neural Information Processing Systems, Cambridge, MA: MIT Press, vol. 8, 1996. – p. 197-203
- 35 Larose D. T. Discovering Knowledge in Data: An Introduction to Data Mining / D. T. Larose. – Wiley & Sons, Inc, 2004. – 240 p.
- 36 Larsen P. M. Industrial applications of fuzzy logic control / P. M. Larsen. // Int. J. Man. Mach. Studies 12, 1980. – Pp. 3-10.
- 37 Loh W. Y. Split selection methods for classification trees / W. Y. Loh, Y. S. Shih // Stat Sin 7, 1997. – Pp. 815–840
- 38 Macintyre A. J. Fitness results for sigmoidal ‘neuronal’ networks / A. J. Macintyre, E. D. Sontag. // Proceedings of 25th Annual ACM Symposium on Theory of Computing. New York: ACM Press, 1993. – Pp. 325-334
- 39 Mamdani E. H. Application of fuzzy algorithms for the control of a simple dynamic plant / E. H. Mamdani // In Proc IEEE, 1974. – Pp. 121–159.
- 40 Mandelbrot B. B. The variation of certain speculative prices / B. B. Mandelbrot // Journal of Business. 1963. No. 36., 1963. – Pp. 394–419.
- 41 Mann L. GOFER : basic principles of decision making / Leon Mann, Ross Harmoni, Colin Power. – Canberra: Curriculum Development Centre, 1988. – 224 p.
- 42 Marwaha S. Analysis and Comparison of Decision Making Ability after Consummation of Customized Solutions and Training Programme / S. Marwaha, G. Marwaha // IOSR Journal of Research & Method in Education, Volume 6, Issue 2 Ver. I (Mar. - Apr. 2016). – Pp. 115–122.
- 43 McCarthy John. What is artificial intelligence? / John. McCarthy – [Electronic resource]. – Available at: \www/URL: – <http://www-formal.stanford.edu/jmc/whatisai/whatisai.html>
- 44 McCulloch W. S. A logical calculus of the ideas immanent in nervous activity / W. S. McCulloch, W. Pitts // Bulletin of Mathematical Biophysics. – 1943. – Vol. 5, 1943. – Pp. 115–133.
- 45 Minsky M. L. Perceptrons / M. L. Minsky, S. A. Papert – Cambridge, MA: MIT Press, 1969. – 263 p.
- 46 Mitchell T. Machine learning / T Mitchell – McGraw-Hill, 1997. – 414 p
- 47 Murphy K. P. Machine Learning: A Probabilistic Perspective (Adaptive Computation and Machine Learning series) / Kevin P. Murphy – 1st Edition. Massachusetts Institute of Technology, 2012. – 1067 p.
- 48 Nash J. F. Non-Cooperative Games / J. F. Nash // Annals of Mathematics 54, 1951. – Pp. 286-295.
- 49 Neelamadhab P. The Survey of Data Mining Applications and Feature Scope / Padhy Neelamadhab , Dr. Mishra Pragnyaban, Rasmita Panigrahi. // International Journal of Computer Science, Engineering and Information Technology, Vol.2, No.3, June 2012. – Pp. 43-58.
- 50 OECD. Understanding the Digital Divide. – Paris. Online-Quelle (Zugriff am 08.11.2002) [Electronic resource]. – Available at: \www/URL:

<http://www.oecd.org/dataoecd/38/57/1888451.pdf>

- 51 Palisade Corporation. Official website. [Electronic resource]. – Available at: \www/URL: <http://www.palisade.com/>
- 52 Power D. J. A Brief History of Decision Support Systems / D. J. Power – DSSResources.COM, World Wide Web, version 4.0, March 10, 2007. [Electronic resource]. – Available at: \www/URL: <http://dssresources.com/history/dsshhistory.html>
- 53 Power D. J., Decision Support Systems: Concepts and Resources for Managers / D. J. Power – Westport, CT: Greenwood/Quorum Books, 2002. – 261p.
- 54 Powers D. M. W. Evaluation: From Precision, Recall and F-Measure to ROC, Informedness, Markedness & Correlation / David M. W. Powers // Journal of Machine Learning Technologies. 2 (1), 2011. – Pp. 37–63.
- 55 Quinlan Ross J. C4.5: Programs for Machine Learning / Ross J. Quinlan // Machine Learning, September 1994, Volume 16, Issue 3, pp 235-240
- 56 Quinlan Ross J. Induction of Decision Trees / Ross J. Quinlan // Machine Learning, Vol. 1, Issue 1 (Mar. 1986), pp. 81-106
- 58 Ramcharan Rodney. Regressions: Why Are Economists Obsessed with Them? / Rodney Ramcharan // International Monetary Fund, Finance & Development, March 2006, Volume 43, Number 1 [Electronic resource]. – Available at: \www/URL: <http://www.imf.org/external/pubs/ft/fandd/2006/03/basics.htm>
- 57 Reshef D. N. Detecting Novel Associations in Large Data Sets / D. N. Reshef, Y. A. Reshef, H K. Finucane and others // Science. № 334 (6062), 2011. – Pp. 1518–1524.
- 59 Rosenblatt F. The perceptron: A probabilistic model for information storage and organization in the brain / F. Rosenblatt // Psychological Review, № 65, 1958. – Pp. 386–407.
- 60 Rumelhart D. E. Learning internal representations by error propagation / D. E. Rumelhart, G. E. Hinton, R. J. Williams in: J. L. McClelland and D. E. Rumelhart (Eds.). // Parallel Distributed Processing: Explorations in the Microstructure of Cognition. Volume 1, MIT Press, Cambridge MA, 1986. – Pp. 318-362.
- 61 Rutkowski L. Metody i techniki sztucznej inteligencji. / L. Rutkowski. – Warszawa.: wydawnictwo naukowe PWN. – 2009. – 452 s.
- 62 Sangameshwari B. Survey on Data Mining Techniques In Business Intelligence / B. Sangameshwari, P. A. Uma // International Journal Of Engineering And Computer Science Volume 3 Issue 10 – October, 2014. – Pp. 8575-8582
- 63 Schauerhuber M. Benchmarking Open-Source Tree Learners in R/RWeka / M. Schauerhuber, A. Zeileis, D. Meyer In C. Preisach, H. Burkhardt, L. Schmidt-Thieme, R. Decker (eds.) // Data Analysis, Machine Learning and Applications (Proceedings of the 31st Annual Conference of the Gesellschaft für Klassifikation e.V., Albert-Ludwigs-Universität Freiburg, March 7–9, 2007). – Pp. 389–396.
- 64 Seddawy Ahmed Bahgat El. Enhanced K-mean Algorithm to Improve Decision Support System under Uncertain Situations / Ahmed Bahgat El Seddawy, Sultan Turkey, Khedr Ayman. // IJCSNS International Journal of Computer

- Science and Network Security, vol.13 No.7, July 2013. – Pp. 50-58.
- 65 Sugeno M. Successive Identification of a Fuzzy Model and its Application to Predict of a Complex System / M. Sugeno, K. Tanaka // Fuzzy Sets Systems, Vol.42, no.3, 1991. – Pp. 315-334.
- 66 Sugeno M., Industrial applications of fuzzy control / M. Sugeno. – Elsevier Science Pub. Co., 1985. – 269 p.
- 67 Tsukamoto Y. Failure Diagnosis by Using of Fuzzy Logic / Y. Tsukamoto, T. Terano, IEEE Proc. Decision & Control, vol. 2., New Orleans, 1977. – Pp. 1390-1395.
- 68 Turban E. Decision Support Systems and Intelligent Systems / Efraim Turban, Jay E. Aronson, Liang Ting-Peng – NJ: Prentice Hall, 2008. – 574 p.
- 69 Vapnik V. N. On the uniform convergence of relative frequencies of events to their probabilities / V. N. Vapnik, A. Ya. Chervonenkis // Theoretical Probability and Its Applications, vol. 17, 1971. – Pp. 264-280.
- 70 Zadeh L. Fuzzy Sets / L. Zadeh // Information and Control, № 8, 1965. – Pp. 338–353.
- 71 Zwicky F. Discovery, Invention, Research through the morphological approach / F. Zwicky. – McMillan, NY, 1969. – 276 p.
- 72 Аверкин А. Н. Толковый словарь по искусственному интеллекту / А. Н. Аверкин, М. Г. Гаазе-Рапопорт, Д. А. Поспелов. – М.: Радио и связь, 1992. – 256 с.
- 73 Азимов Э. Г. Новый словарь методических терминов и понятий (теория и практика обучения языкам) / Э. Г. Азимов, А. Н. Щукин. – М.: Издательство ИКАР, 2009. – 448 с.
- 74 Акимов С. В. Анализ проблемы автоматизации структурно-параметрического синтеза / С. В. Акимов // Доклады ТУСУР. – Томск, 2011. – №2-2 (24), 2011. – С.204-211.
- 75 Акимов С. В. Два способа задания множества альтернатив при формализации задачи структурно-параметрического синтеза / С. В. Акимов // Методы и алгоритмы прикладной математики в технике, медицине и экономике: материалы VI Международной научно практической конференции. Ч. 1 / Новочеркасск, 2006. – С. 11–12
- 76 Акимов П.С. Сигналы и их обработка в информационных системах / П. С. Акимов, А. И. Сенин, В. И. Соленов. – М.: Радио и связь, 1994. – 296 с.
- 77 Алферова З. В. Математическое обеспечение экономических расчетов с использованием теории графов / З. В. Алферова. – М.: Статистика, 1974. – 208 с.
- 78 Альтшуллер Г. С. Найти идею: Введение в ТРИЗ – теорию решения изобретательских задач / Г. С. Альтшуллер. – 4-е изд. – М.: Альпина Паблишерз, 2011. – 400 с.
- 79 Альтшуллер Г. С. Творчество как точная наука / Г. С. Альтшуллер. – М.: Сов. радио, 1979. – 175 с.
- 80 Аналітична довідка «Стан розвитку науки і техніки, результати наукової, науково-технічної, інноваційної діяльності, трансферу технологій за 2015

рік» / Український інститут науково-технічної і економічної інформації [Електронний ресурс] – Режим доступу: <http://mon.gov.ua/content/Діяльність/Наука/2-3-ad-kmu-2015.pdf> – 199 с.

- 81 Анфилов В. С. Системный анализ в управлении / В. С. Анфилов, А. А. Емельянов, А. А. Кукушкин. – М.: Финансы и статистика, 2002. – 368 с.
- 82 Афанасьев В. Г. Общество: системность, познание, управление / В. Г. Афанасьев. – М.: Политиздат, 1981. – 432 с.
83. Батищев Д. И. Сравнение эффективности популяционно-генетического и классического подходов для решения задачи адаптации оптимальных решений нестационарной задачи комбинаторной оптимизации / Д. И. Батищев, Е. А. Неймарк // Вестник Нижегородского университета им. Н. И. Лобачевского. – Н.Новгород: НУ им. Лобачева. –2009. – Вып №5. – С. 169-172.
- 84 Батракова Л. Г. Теория статистики / Л. Г. Батракова – М.: КноРус, 2013. – 528 с.
- 85 Белоусов А. И. Дискретная математика / А. И. Белоусов, С. Б. Ткачев. – М.: МГТУ, 2004. – 744 с.
- 86 Берндт Э. Р. Практика эконометрики классика и современность. Пер. с англ. / Э. Р. Берндт. – М. ЮНИТИ-ДАНА. 2005. – 863 с.
- 87 Бир С. Мозг фирмы / С. Бир. – М.: Едиториал УРСС, 2005. – 416 с.
- 88 Бирштейн Б. И. Стратегемы рефлексивного управления в западной и восточной культурах / Б. И. Бирштейн, В. И. Боршевич // Рефлексивные процессы и управление. – 2002. – Т.2. - №1. – С.27-44.
- 89 Божко А. Н. Структурный синтез как задача дискретной оптимизации / А. Н. Божко //Электронное научно-техническое издание «Наука и Образование», 2010, №9. [Электронный ресурс] – Режим доступа: <http://technomag.edu.ru/doc/158337.html>
- 90 Божко А. Н., Толпаров А. Ч. Структурный синтез на элементах с ограниченной сочетаемостью / А. Н. Божко // Электронное научно-техническое издание «Наука и Образование», 2004, №5. [Электронный ресурс] – Режим доступа: <http://www.metodolog.ru/00562/00562.html>
- 91 Большой словарь цитат и крылатых выражений // сост. К. Душенко/ М.: Эксмо, ИНИОН РАН, 2011. – 1216 с.
- 92 Боресков А. В. Основы работы с технологией CUDA / А. В. Боресков, А. А. Харламов. – ДМК-Пресс, 2010. – 232 с.
- 93 Боровский В. Н. Методы и проблемы анализа финансовой устойчивости банков в Украине / В. Н. Боровский, Я. А. Гатинский // Культура народов Причерноморья. – № 201, 2011. – С. 15-18
- 94 Буркун И. Г. Формирование цены предложения на рынке жилой недвижимости региона / И. Г. Буркун // Інвестиції: практика та досвід. – 2010. – № 7. – С. 49-52.
- 95 Бюлетень Національного Банку України. [Електронний ресурс]. – Режим доступу: <http://www.bank.gov.ua/>
- 96 Ватулин Э. И. Основы дискретной комбинаторной оптимизации / Э. И. Ватулин, В. С. Титов, С. Г. Емельянов. – М.: АРГАМАК-МЕДИА,

2016. – 270 с.

- 97 Використання інформаційно-комунікаційних технологій на підприємствах України Статистичний бюлетень / Державна служба статистики України, Київ, 2015 [Електронний ресурс] - Режим доступу: http://www.ukrstat.gov.ua/druk/publicat/kat_u/2015/bl/07/bl_vikt_14.zip
- 98 Висновки щодо виконання закону про державний бюджет України на 2015 рік [Електронний ресурс]. – Режим доступу: http://www.ac-rada.gov.ua/doccatalog/document/16748378/Vykonan_DBU_2015.pdf
- 99 Волков В. А. Системный анализ для структурно-параметрического синтеза / В. А. Волков, С. М. Чудинов // Научные ведомости Белгородского государственного университета. Серия: Экономика. Информатика. – Белгород, 2012. – № 19-1 (138) / том 24. – С. 153-157.
- 100 Волкова В. Н. Системный анализ и принятие решений: Словарь-справочник / Под ред. В. Н. Волковой, В. Н. Козлова. – М.: Высш. шк., 2004. – 616 с.
- 101 Вьюгин В. Математические основы машинного обучения и прогнозирования / В. Вьюгин. – М.: МЦМНО, 2014. – 304 с.
- 102 Газизов Д. И. Обзор методов статистического анализа временных рядов и проблемы, возникающие при анализе нестационарных временных рядов /Д. И. Газизов // Научный журнал. – 2016. – №3. – С. 9-14
- 103 Гладков Л. А. Генетические алгоритмы / Л. А. Гладков, В. В. Курейчик, В. М. Курейчик – М.: ФИЗМАТЛИТ, 2006. – 320 с.
- 104 Губко М. В. Классификация моделей анализа и синтеза организационных структур / М. В. Губко, Н. А. Коргин // Управление большими системами: сборник трудов. – М., 2004 – №6. – С. 5-21.
- 105 Декарт Р. Рассуждение о методе, чтобы верно направлять свой разум и отыскать истину в науках. Метафизические размышления. / Р. Декарт. – Начала философии. – М.: Вежа, 1998. – 240 с.
- 106 Державний комітет статистики України. [Електронний ресурс]. – Режим доступу: <http://ukrstat.gov.ua>
- 107 Джарратано Дж. Экспертные системы: принципы разработки и программирование: Пер. с англ. / Джозеф Джарратано, Гарри Райли. – М.: Издательский дом «Вильямс», 2006. – 1152 с.
- 108 Дорошенко А. Ю. Формалізоване проектування та синтез паралельних програм для відеографічних прискорювачів /А. Ю. Дорошенко, О. Г. Бекетов, К. А. Жереб, О. А. Яценко // Проблеми програмування. 2013. № 3. – С. 38 – 46.
- 109 Дьяконов В. Математические пакеты расширения Matlab. Специальный справочник / В. Дьяконов, В. Круглов. – СПб.: Питер, 2001. – 480 с
- 110 Ежов А. А. Нейрокомпьютинг и его применения в экономике и бизнесе / А. А. Ежов, С. А. Шумский., под ред. проф. В. В. Харитоновой – серия «Учебники экономико-аналитического института МИФИ». М.: МИФИ, 1998. – 224 с.
- 111 Економіка України: Економічні новини. [Електронний ресурс]. – Режим доступу: <http://www.ereport.ru/articles/weconomy/ukraine.htm>
- 112 Жук О. В. Стан і перспективи розвитку іпотечного кредитування в

- Україні / О. В. Жук // Економічний простір. – 2009. – № 23/1. – С. 308-315.
- 113 Заде Л. Понятие лингвистической переменной и ее применение к принятию приближенных решений: Пер. с англ. / Л. Заде – М.: Мир, 1976. – 167 с.
- 114 Зак Ю. А. Принятие решений в условиях нечетких и размытых данных: Fuzzy-технологии / Ю. А. Зак – М.: Книжный дом "Либроком", 2013. – 352 с.
- 115 Интерактивный график EUR/USD. [Электронный ресурс] - Режим доступа: <http://ru.investing.com/currencies/eur-usd-advanced-chart>
- 116 Истомин Л. Ф. Логические основы систем управления / Л. Ф. Истомин, В. К. Зайко, С. М. Танченко. – Луганск: ВНУ им. В. Даля, 2005. – 335с.
- 117 Ишина И. В. Скоринг-модель оценки кредитного риска / И. В. Ишина, М. Н. Сазонова // Аудит и финансовый анализ. – 2007. – № 4. – с. 297-304.
- 118 Івасів І. Б. Актуальні підходи до банківського моніторингу / І. Б. Івасів // Фінанси, облік і аудит : зб. наук. праць – К.: КНЕУ, 2004. – Вип. № 4. – С. 86–94.
- 119 Каменський А. Б. Економіко-математичне моделювання фінансових ризиків: Автореф. дис. докт. екон. наук: 08.00.11 / А. Б. Каменський. – КНУ ім. Тараса Шевченка. – К., 2007. – 34 с.
- 120 Канторович Л. В. Математические методы организации и планирования производства / Л. В. Канторович. – Л. Изд-во ЛГУ, 1939. – 68 с.
- 121 Карпенко А. П. Современные алгоритмы поисковой оптимизации. Алгоритмы, вдохновленные природой/ А. П. Карпенко. – М: Издательство МГТУ им. Н. Э. Баумана, 2014. – 446с.
- 122 Карпов Ю. Имитационное моделирование систем. Введение в моделирование с AnyLogic 5 / Ю. Карпов. – СПб.: БХВ-Петербург, 2005. – 400 с.
- 123 Кнорринг В. И. Теория, практика и искусство управления / В. И. Кнорринг. – М.: НОРМА, 2001. – 528 с.
- 124 Кнут Д. Искусство программирования, том 1. Основные алгоритмы = The Art of Computer Programming, vol.1. Fundamental Algorithms /Дональд Кнут – М.: Вильямс, 2006. – 720 с.
- 125 Ковалев В. В. Анализ хозяйственной деятельности предприятия / В. В. Ковалев, О. Н. Волкова. – М.: ТК Велби., 2002. – 424 с.
- 126 Корилов А. М. Теория систем и системный анализ / А. М. Корилов, С. Н. Павлов – Томск: Томский гос. ун-т систем управления и радиоэлектроники, 2008. – 264 с.
- 127 Корсаков С. Н. Начертание нового способа исследования при помощи машин, сравнивающих идеи. Пер. с франц. / С. Н. Корсаков, под ред. А. С. Михайлова. – М.: МИФИ, 2009, 44 с.
- 128 Коэлья Л. П., Построение систем машинного обучения на языке Python / Л. П. Коэлья, В. Ричард. – М.: ДМК Пресс, 2016. – 302 с.
- 129 Кравчук Е. В. Искусственные нейронные сети и генетические алгоритмы /Е. В. Кравчук., Э. Хантер. – Донецк: ДонГУ, 2000. – 200 с.
- 130 Ларичев О. И. Системы поддержки принятия решений: современное состояние и перспективы развития / О. И. Ларичев, А. Б. Петровский. – Итоги науки и техники. М.: ВИНТИ. – 1987. –Т. 21. – С. 131-164.

- 131 Лернер Ю. И. Оценка финансовой устойчивости банковской структуры / Ю. И. Лернер // Вісник економічної науки України. – 2011. – № 2 (20). – С. 82-86.
- 132 Лефевр В. А. Алгебра Совести / В. А. Лефевр. – М.: Когито-центр, 2003. – 426 с.
- 133 Лефевр В. А. Просчеты миротворчества / В. А. Лефевр // Рефлексивные процессы и управление, 2002. – Т.2. – №2.– С.48-52.
- 134 Лефевр В. А. Рефлексия / В. А. Лефевр. – М., «Когито-Центр», 2003. – 496 с.
- 135 Лисенко Ю. Г. Управление коммерческим банком: инновационный аспект / Ю. Г. Лысенко, В. Н. Тимохин, Р. А. Руденький и др. // Донецк.: ООО «Юго-Восток, Лтд. 2008. – 328 с.
- 136 Лопатников Л. И. Экономико-математический словарь: словарь современной экономической науки / Л.И. Лопатников. – М.: Дело, 2003.– 520 с.
- 137 Лысенко Ю. Г. Поиск эффективных решений в экономических задачах / Ю. Г. Лысенко, А. Ю. Минц, В. Г. Стасюк, – Донецк: ДонНУ; ООО «Юго-Восток, Лтд», 2002. – 101 с.
- 138 Лысенко Ю. Г., Иванов Н. Н., Минц А. Ю. Нейронные сети и генетические алгоритмы / Ю. Г. Лысенко, Н. Н. Иванов, А. Ю. Минц.– Донецк: ДонНУ, ООО «Юго-Восток, Лтд», 2003. – 230 с.
- 139 Макаров И. М. Теория выбора и принятия решений / Макаров И. М., Виноградская Т. М., Рубчинский А. А., Соколов В. Б. // М.: Наука, 1982. – 328 с.
- 140 Марков А. А. Теория алгорифмов / А. А. Марков, Н. М. Нагорный – М.: Наука, 1984. – 432 с.
- 141 Мастицкий С. Э. Статистический анализ и визуализация данных с помощью R / С. Э. Мастицкий, В. К. Шитиков. – М.: ДМК Пресс, 2015. – 496 с.
- 142 Масютина Г. В. Структурно-параметрический синтез адаптивной системы управления на основе нечеткой логики / Г. В. Масютина, В. Ф. Лубенцов // Известия Южного федерального университета. – Ростов-на-Дону. – 2010. – № 5, т.106. – С. 165-170
- 143 Матвійчик А. В. Штучний інтелект в економіці: нейронні мережі, нечітка логіка: монографія / А. В. Матвійчук. – К. : КНЕУ, 2011. – 439 с.
- 144 Мертенс А. Инвестиции. /А. Мертенс. – Киев: Киевское инвестиционное агентство, 1997. - 347 с.
- 145 Минеев А. А. Разработка инструментария планирования социально-ориентированного развития экономики промышленных предприятий: автореф. дис. на соиск. учен. степ. канд. эконом. наук / А. А. Минеев. – Московский фин-юр. университет МФЮА., Москва, 2012. – 24с.
- 146 Минц А. Ю. Анализ интеллектуальных средств поддержки принятия решений в экономических задачах [Электронный ресурс] / А. Ю. Минц // Економіка та суспільство.: Електронне наукове фахове видання. – Мукачево, 2016. – №2/2016. – С. 784-790. – Режим доступу: <http://economyandsociety.in.ua>

- 147 Минц А. Ю. Генетические алгоритмы оптимизации рефлексивных воздействий. / А. Ю. Минц, Е. В. Хаджинова, М. И. Никонова // Вісник Приазовського державного технічного університету. Сер.: Економічні науки: Зб. наук. праць. – Маріуполь: ДВНЗ «ПДТУ», 2012. – Вип. 24. – С. 75-84.
- 148 Минц А. Ю. Инструментальные средства генетического моделирования и перспективы их использования для поиска оптимальных решений экономических задач./ А. Ю. Минц // – Нове в економічній кібернетиці: зб. наук. ст.. / під загал. ред.. Ю. Г. Лисенко; Донецький нац. ун-т. – Донецьк: «Юго-Восток», 2010. – Вип. 4: Технології штучних нейронних мереж в економіці. – С. 79-96
- 149 Минц А. Ю. Интеллектуальные методы анализа надежности участников рынков финансовых услуг. / А. Ю. Минц // Вісник Донецького університету економіки та права. : зб. наук. праць. – Артемівськ: ДонУЕП, 2015. – № 2/2015. – С. 85-90.
- 150 Минц А. Ю. Концептуальные подходы к моделированию интеллектуальных автоматизированных систем принятия решений. / А. Ю. Минц // Нове в економічній кібернетиці : зб. наук. ст. під загал. ред. Ю. Г. Лисенко; Донецький нац. ун-т. – Донецьк, 2014. – вип 3/2014. – С. 70-81.
- 151 Минц А. Ю. Концепция моделирования интеллектуальных автоматизированных систем принятия решений в управлении экономическими объектами. / А. Ю. Минц // Вісник Донецького національного університету. Серія В «Економіка і право». – Вінниця, 2015. – №1/2015. – С.253-258.
- 152 Минц А. Ю. Краудсорсинг, как метод решения задач в глобализованной экономике и особенности его использования в Украине / А. Ю. Минц // Вісник Приазовського державного технічного університету : зб. наукових праць / ПДТУ. – Маріуполь, 2013. – Вип. 26 (Серія : Економічні науки). – С. 85-90.
- 153 Минц А. Ю. Львовский Л. Я. Нейро-скоринговый метод оценки кредитоспособности заемщиков / А. Ю. Минц, Л. Я. Львовский // Нове в економічній кібернетиці: зб. наук. ст. під загал. ред. Ю. Г. Лисенко; Донецький нац. ун-т. – Донецьк: «Юго-Восток», 2010. – Вип. 4: Технології штучних нейронних мереж в економіці. – С. 70-79.
- 154 Минц А. Ю. Метод определения доходов различных групп населения / А. Ю. Минц // Научный взгляд в будущее. – Иваново: ООО "Научный мир", 2016. – Выпуск 2(2). Том 7. – с. 71-75
- 155 Минц А. Ю. Метод упрощения динамических рядов с использованием генетических алгоритмов / А. Ю. Минц // Економічний вісник запорізької державної інженерної академії. – Запоріжжє, 2016. – Вип. 4(04) Часть 2. – С. 120-124
- 156 Минц А. Ю. Методы отбора данных для нейросетевого моделирования / А. Ю. Минц // Моделювання та інформаційні системи в економіці, зб.наук.пр. – Київ: КНЕУ, 2011. – Вип. 84. – С. 256-270.

- 157 Минц А. Ю. Методы оценки эффективности решения задач ранжирования / А. Ю. Минц // Економічна кібернетика: міжн. науч. журнал / ДонНУ. – Донецьк: Юго-Восток, Лтд, 2012. – №1-3(73-75). – С. 51-56
- 158 Минц А. Ю. Моделирование процессов продвижения услуг электронного бизнеса коммерческих банков / А. Ю. Минц // Новое в экономической кибернетике (сб. н. ст.); Донецкий нац. ун-т.// Национальная экономика: методы, модели, механизмы. – Донецк: Юго-Восток, 2009. – №3. – с. 219-228
- 159 Минц А. Ю. Моделирование финансового состояния заемщиков – физических лиц в кризисных условиях. / А. Ю. Минц, Л. Я. Львовский // Вісник Запорізького національного університету. – 2010. – №4(8). – С. 117-123
- 160 Минц А. Ю. Моделирование ценообразования на рынке жилой недвижимости методами системной динамики / А. Ю. Минц // Технологический аудит и резервы производства. – Харьков, 2016. – №5/4(31). – С. 39- 45.
- 161 Минц А. Ю. Общие вопросы постановки задач в нейросетевом моделировании / А. Ю. Минц // Нейро-нечіткі технології моделювання в економіці, наук.-аналіт. журн. – Київ: КНЕУ, 2012. – №1. – С. 189-206.
162. Минц А. Ю. Оптимизация затратной части инновационных проектов / А. Ю. Минц // Теоретичні та практичні аспекти економіки та інтелектуальної власності. – Збірник наукових праць. – Маріуполь: ПДТУ, 2008. – с 116-119.
- 163 Минц А. Ю. Прогнозирование валютных рынков с использованием самоорганизующихся нейронных сетей / А. Ю. Минц // Вісник СНУ ім. В.Даля. – 2004. – №4(74). – с. 184-193.
- 164 Минц А. Ю. Управление кредитным рейтингом как способ оптимизации расходов заёмщика /А. Ю. Минц, С. С. Глушаченко // Науковий вісник Херсонського державного університету. Серія «Економічні науки». – Херсон, 2014. – Вип. 8, част.5. – С.173-176
- 165 Минц А. Ю. Формализованный метод анализа конкурентоспособности банковских платежных карт. / А. Ю. Минц // Вісник ПДТУ. сер.: Економічні науки: зб. наук. праць. – Маріуполь: ПДТУ, 2010. – Вип. 20. – С. 96-101.
- 166 Минц А. Ю. Генетическая модель оптимизации рефлексивных воздействий при взаимодействии предприятия с потребителями / А. Ю. Минц, Е. Л. Петрачкова // Вісник економічної науки України. – 2006. – №2(10). – С. 129-134
- 167 Минц А. Ю., Современные методы анализа данных в финансово-кредитной сфере / А. Ю. Минц, Е. В. Хаджинова // Вісник Приазовського державного технічного університету. Сер.: Економічні науки: Зб. наук. праць. - Маріуполь: ДВНЗ «ПДТУ», 2011. – No.2 (22). – С. 149-156.
- 168 Мир на рубеже тысячелетий (прогноз развития мировой экономики до 2015 г.) / Руководители авторского коллектива академик РАН В.А. Мартынов, член-корр. РАН А.А. Дынкин. – М., Издательский Дом НОВЫЙ ВЕК, 2001.

– 592 с.

- 169 Миронов А. М. Теория процессов [Электронный ресурс]./ А. М. Миронов. – Переславль-Залесский: «Университет города Переславля», 2008. – 345с. – Режим доступа: <http://is.ifmo.ru/verification/mironov-process-theory.pdf>
- 170 Мінц О. Ю. Інвестиційний механізм у фінансовій санації підприємств / О. Ю. Мінц, О. В. Хаджинова // Теоретичні і практичні аспекти економіки та інтелектуальної власності.: Зб. наук. праць. – Маріуполь: ДВНЗ «ПДТУ», 2012. – Вип. 1, Т.3. – С.127–131.
- 171 Мінц О. Ю. Інтелектуальні методи прогнозування рівня виконання державного бюджету України [Електронний ресурс]/ О. Ю. Мінц // Глобальні та національні проблеми економіки.: Електронне наукове видання. – Миколаїв, 2016. – №12/2016. – С. 573-580. –Режим доступу: <http://www.global-national.in.ua/archive/12-2016/118.pdf>
- 172 Мінц О. Ю. Методи прогнозування кількості банкрутств в Україні./ О. Ю. Мінц, К. Е. Беззубкова // Економіка і організація управління. – Вінниця, 2014. – № 1 (17) - 2 (18). – С. 172-179
- 173 Мінц О. Ю. Методы анализа конкурентоспособности коммерческих банков в сфере электронного бизнеса. / О. Ю. Мінц, Л. С. Омельченко О. В. Хаджинова // Збірник наукових праць Донецького державного університету управління «Маркетинг і економічний розвиток підприємств», серія “Економіка”. – Донецьк: ДонДУУ, 2009. – т. 10, вип. 144. – С.169-176.
- 174 Мінц О. Ю. Моделювання процесів реструктуризації кредитів // Вісник Університету банківської справи НБУ: Зб. наук. праць. – Київ: УБС НБУ, 2012. – №2(14). - С. 329-333
- 175 Мінц О. Ю. Формування інноваційної стратегії в системі антикризового управління підприємством./ О. Ю. Мінц, О. В. Хаджинова // Теоретичні і практичні аспекти економіки та інтелектуальної власності: Збірник наукових праць. – Маріуполь: ПДТУ, 2010. – Т.2. – С.206-210.
- 176 Міщенко В. Реструктуризація кредитів в умовах кризи: світовий досвід і можливості застосування в Україні / В. Міщенко, В. Крилова, М. Ніконова // Вісник НБУ. – 2009. – № 5. – С. 12–17.
- 177 Неітеративні, еволюційні та мультиагентні методи синтезу нечіткологічних і нейромережних моделей: Монографія / Під заг. ред. С. О. Субботіна. – Запоріжжя: ЗНТУ, 2009. – 375 с.
- 178 Нейман Дж. Теория игр и экономическое поведение / Дж. Нейман, О. Моргенштерн. – М.: Наука, 1970. – 708 с.
- 179 Неплатоспроможні банки. Фінансовий портал Мінфін. [Електронний ресурс]. – Режим доступу: <http://minfin.com.ua/banks/problem/>
- 180 Нечеткие модели и нейронные сети в анализе и управлении экономическими объектами: монография / под ред. Ю. Г. Лысенко. – Донецк: Юго-Восток, 2012. – 388с.
- 181 Новиков А. М. Методология / А. М. Новиков, Д. А. Новиков. – М.: СИНТЕГ, 2007. – 668 с.
- 182 Норенков И. П. Основы автоматизированного проектирования /

- И. П. Норенков. – М.: Изд-во МГТУ им. Н.Э. Баумана, 2009. – 430 с.
- 183 Сайт компанії Mathworks, виробника програмного продукту MatLab [Електронний ресурс] – Режим доступу: <https://www.mathworks.com/products/matlab.html>
- 184 Сайт системи виконання математичних розрахунків Octave [Електронний ресурс] – Режим доступу: <http://www.gnu.org/software/octave/>
- 185 Сайт компанії StatSoft [Електронний ресурс] – Режим доступу: <http://www.statsoft.ru/>
- 186 Сайт компанії NeuroDimension, виробника програмного продукту Neuro Solutions [Електронний ресурс] – Режим доступу: <http://www.neurosolutions.com/>
- 187 Сайт програмного продукту NeuroShell Trader [Електронний ресурс] – Режим доступу: <http://try.neuroshell.com/index/>
- 188 Сайт компанії, виробника програмного продукту Neural Designer [Електронний ресурс] – Режим доступу: <http://neuraldesigner.com>
- 189 Сайт компанії BaseGroup, виробника програмного продукту Deductor Studio [Електронний ресурс] – Режим доступу: <http://www.basegroup.ru/>
- 190 Сайт компанії Palisade Corporation, виробника програмного продукту Evolver [Електронний ресурс] – Режим доступу: <http://www.palisade.com/>
- 191 Сайт компанії Ward Systems [Електронний ресурс] – Режим доступу: <http://www.wardsystems.com/index.asp>
- 192 Огляд банківського сектору, лютий 2017 р. [Електронний ресурс] / Офіційний сайт НБУ. – Режим доступу: <https://bank.gov.ua/doccatalog/document?id=43633516>
- 193 Орлов А. И. Математические методы теории классификации [Электронный ресурс]/ А. И. Орлов // Политематический сетевой электронный научный журнал Кубанского государственного аграрного университета (Научный журнал КубГАУ). – Краснодар: КубГАУ, 2014. – №01(095). – С. 423–459. Режим доступа: <http://ej.kubagro.ru/2014/01/pdf/23.pdf>
- 194 Орлюк О. П. Фінансове право / О. П. Орлюк. - К.: Юрінком Інтер, 2003. - 527 с.
- 195 Паклин Н. Б., Бизнес-аналитика: от данных к знаниям / Н. Б. Паклин, В. И. Орешков. – СПб.: Питер, 2013. – 704 с.
- 196 Панченко Т. В. Генетические алгоритмы / Т. В. Панченко, под ред. Ю. Ю. Тарасевича. – Астрахань: Издательский дом "Астраханский университет", 2007. – 87 с.
- 197 Петерс Э. Фрактальный анализ финансовых рынков: Применение теории Хаоса в инвестициях и экономике / Э. Петерс. – М.: Интернет-Трейдинг, 2004. – 304 с.
- 198 Петерс Э. Хаос и порядок на рынках капитала. Новый аналитический взгляд на циклы, цены и изменчивость рынка / Э. Петерс. – М.: Мир, 2000. – 333 с.
- 199 Положення про порядок формування та використання резерву для відшкодування можливих втрат за кредитними операціями банків: Постанова правління НБУ №279 від 06.07.2000. із змінами та

- доповненнями [Електронний ресурс] – Режим доступу: <http://zakon2.rada.gov.ua/laws/show/z0231-12>
- 200 Проникновение Интернета в Украине на 1-й квартал 2017 года / Factum Group Ukraine, 2017 [Електронний ресурс] - Режим доступу: <https://www.slideshare.net/WatcherUA/internet-audience-in-ukraine-1q-2017>
- 201 Пятецкий-Шапиро Г., Data Mining и перегрузка информацией // Вступительная статья к книге: Анализ данных и процессов / А. А. Барсегян, М. С. Куприянов, И. И. Холод, М. Д. Тесс, С. И. Елизаров. 3-е изд. перераб. и доп. СПб.: БХВ-Петербург, 2009. – 512 с., с.13.
- 202 Радіонов Ю. Д. Прогнозування і планування як інструмент ефективного управління та використання бюджетних коштів / Ю. Д. Радіонов // Економіка України. – 2014. – № 4. – С. 40-54.
- 203 Рашкован В. Кластерний аналіз бізнес-моделей українських банків: застосування нейронних мереж Кохонена / В. Рашкован, Д. Покідін // Вісник НБУ. – 2016. – № 238. – С 13-40.
- 204 Рогозин В. О. Сравнительная оценка алгоритмов поддержки принятия решений на основе качественных характеристик. / В. О. Рогозин // Образовательные технологии, 2010. – № 4/2010. – С. 84-105.
- 205 Росс Эшби У. Введение в кибернетику / Эшби Уильям Росс. – М.: Издательство иностранной литературы, 1959. – 432 с.
- 206 Румянцева З. П. Общее управление организацией. Теория и практика/ З. П. Румянцева. – М.: ИНФРА – М, 2003. – 304 с.
- 207 Сараев А. Д. Системный анализ и современные информационные технологии / А. Д. Сараев, О. А. Щербина. – Симферополь: СОНАТ, 2006. – 342 с.
- 208 Сарычева Л. В. Кластерно-регрессионный анализ финансовых показателей банков Украины на основе МГУА / Л. В. Сарычева, А. П. Сарычев // Індуктивне моделювання складних систем: Зб. наук. пр. – К.: МННЦ ІТС НАН та МОН України, 2013. – Вип. 5. – С.270-277
- 209 Свами М. Графы, сети и алгоритмы/ М Свами, К. Тхуласираман. – М. "Мир", 1984. – 454 с.
- 210 Седжвик Р. Алгоритмы на С++ / Р. Седжвик. – М.: Вильямс, 2017. – 1056 с.
- 211 Сергеева Л. Н. Нелинейная экономика: модели и методы / Научн. редактор д.э.н., проф. Ю. Г. Лысенко. – Запорожье: Полиграф, 2003. – 218 с.
- 212 Сирота В. Реструктуризація позик, як ефективний інструмент управління проблемними активами / В. Сирота // Вісник Університету банківської справи Національного банку України. – 2011. – № 3 (12). – С.207-210.
- 213 Систематизация // Большая советская энциклопедия : [в 30 т.] / гл. ред. А. М. Прохоров. – 3-е изд. – М. : Советская энциклопедия, 1969–1978.
- 214 Сіташ Т. Д. Планування бюджетних видатків: концептуалізація та тенденції / Т. Д. Сіташ. // Економіка. Управління. Інновації. – 2014. – № 1 (11). – С. 164-169.
- 215 Скворцова Н. А. Информационно-коммуникационная среда вуза как средство формирования профессионализма студентов [Электронный ресурс]/ Н. А. Скворцова, Н. В. Пьянова // Режим доступа:

- http://www.rai.ru/snt/?section = content&op = show_article&article_id = 5651
- 216 Стандарт ISO/IEC TR 9126-4:2004 [Електронний ресурс] Режим доступу: <https://www.iso.org/standard/39752.html>
- 217 Статистика цін на житлову нерухомість у м. Київ [Електронний ресурс]. – Режим доступу: <http://realt.ua/Db2/0st.php?Opr=1>
- 218 Сушков Ю. А. Гипердеревья и блок-схемы механизмов / Ю. А. Сушков // Дискретные модели. Анализ, синтез и оптимизация. – СПбГУ, 1998. – С.40-47.
- 219 Тихомиров В. П. Мир на пути к Smarteducation: новые возможности для развития [Электронный ресурс] / Режим доступа: <http://www.slideshow.net/ssusere58270/e-learning -240511>.
- 220 Томас Т. Л. Рефлексивное управление в России: теория и военные приложения / Т. Л. Томас // Рефлексивные процессы и управление. – 2002. – Т.2. №1. – С.71-89.
- 221 Тоффлер Э. Третья волна = The Third Wave, 1980 / Э. Тоффлер. – М.: АСТ, 2010. – 784 с.
- 222 Уланов С. В. Оценка качества и сравнение скоринговых карт/ С. В. Уланов // Экономические науки. –2009. – № 9(58). – С 330-335.
- 223 Урсул А. Д. Проблема информации в современной науке / А. Д. Урсул. – М.: Наука, 1975. – 288 с.
- 224 Успенский В. А. Теорема Гёделя о неполноте / В. А. Успенский . – М.: Наука, 1982. 110 с.
- 225 Философский словарь. Под ред. М.М. Розенталя. Изд. третье. – М.: Изд-во политической литературы, 1972. – 495 с.
- 226 Философский энциклопедический словарь / Гл. редакция: Л. Ф. Ильичев, П. Н. Федосеев, С. М. Ковалев, В. Г. Панов. – М.: Сов. Энциклопедия, 1983. – 840 с.
- 227 Финансовый словарь [Электронный ресурс] // М.: ФИНАМ – Режим доступа – <http://www.finam.ru/dictionary>
- 228 Фінансова звітність банків України. [Електронний ресурс]. – Режим доступу: http://bank.gov.ua/control/uk/publish/category?cat_id=64097
- 229 Фомін Г. Ф. Фінансово-правове забезпечення касового виконання Державного бюджету / Г. Ф. Фомін // Вісник Харківського національного університету внутрішніх справ. – 2000. – № 10. – С. 199-203.
- 230 Фонд гарантування вкладів фізичних осіб. Інформація щодо виведення банків з ринку [Електронний ресурс]. – Режим доступу: <http://www.fg.gov.ua/not-paying>
- 231 Форрестер Дж. Основы кибернетики предприятия (Индустриальная динамика) / Дж. Форрестер. Пер. с англ. – М.: Прогресс, 1971. – 466 с.
- 232 Хайкин С. Нейронные сети: Полный курс / С. Хайкин – [2-е изд.]. – М.: Вильямс, 2006. – 1104 с.
- 233 Хатимлянський А. Генетическіє алгоритми в MetaTrader 4. Сравнение с прямым перебором оптимизатора. [Электронный ресурс]/ А. Хатимлянський. – Режим доступа: <http://articles.mql4.com/ru/135>
- 234 Хмарук Ю. В. Оцінка факторів впливу на доходи Державного бюджету

- України / Ю. В. Хмарук // Наукові записки Національного університету "Острозька академія". Сер. : Економіка. – 2011. – Вип. 16. – С. 82-91.
- 235 Черняк Л. С. ВІ и DSS - две стороны одной медали / Л. С. Черняк // Открытые системы. СУБД. – Москва: Открытые системы. – 2009. – № 9. – С.22-26.
- 236 Черняк Ю. И. Системный анализ в управлении экономикой / Ю. И. Черняк. – М.: Экономика, 1975. – 193 с.
- 237 Чуркин Г. М. Использование комбинаторно-логических методов структурного синтеза при выборе концепции технического обеспечения автоматизированных систем / Г. М. Чуркин // Вестник СГТУ. – Саратов, 2014. – № 3 (76). – С. 97-104.
- 238 Щедровицкий П. Г. К анализу топики организационно-деятельностных игр. Препринт / П. Г. Щедровицкий. – М.: Научн. центр биол. иссл-й АН СССР в Пущино, 1987. – 43 с.
- 239 Экономика и право. Энциклопедический словарь // М. изд-во Дашков и Ко. 2001. – 568 с.
- 240 Эрлих А. Теория и практика валютного дилинга / А. Эрлих. – К.: 1998. – 180 с.
- 241 Эрлих А. Технический анализ товарных и фондовых рынков / А. Эрлих.– М.: Юнити, 1996. – 318 с
- 242 Юдин Э. Г. Системный подход и принцип деятельности: методологические проблемы современной науки / Э. Г. Юдин. – Москва: Наука, 1978. – 390 с.

ДОДАТОК А

ПРОГРАМНА РЕАЛІЗАЦІЯ НЕЧІТКОЇ МОДЕЛІ ОЦІНКИ ІНСТРУМЕНТАЛЬНИХ ЗАСОБІВ РОЗРОБКИ ІСПР

```
x = (0:0.01:5)';
prod = {'C'      'matlab' 'Deductor' 'Neur Des' 'R' 'Python'};
%задаем оценки привлекательности продукта
y0 = trapmf(x, [0 0 0 0]);
y1 = trapmf(x, [0 0 0.2 0.25]); %очень низкая
y2 = trapmf(x, [0.15 0.25 0.35 0.4]); %низкая
y3 = trapmf(x, [0.3 0.4 0.55 0.6]); %средняя
y4 = trapmf(x, [0.45 0.6 0.7 0.75]); %высокая
y5 = trapmf(x, [0.65 0.75 1 1]); %очень высокая
y=[y1 y2 y3 y4 y5];

%задаем значимость критериев оценки
z1 = trapmf(x, [0 0 0 0.4]); %не значит
z2 = trapmf(x, [0.2 0.35 0.5 0.6]); %средняя значимость
z3 = trapmf(x, [0.45 0.6 0.7 0.8]); %высокая
z4 = trapmf(x, [0.65 0.75 1 1]); %очень высокая

%задаем оценки каждого продукта по 7 критериям
p(1,:)= [1 1 5 4 5 5 5];
p(2,:)= [3 5 3 2 4 3 5];
p(3,:)= [5 3 2 3 4 5 5];
p(4,:)= [4 4 3 2 3 3 3];
p(5,:)= [2 4 5 4 4 3 4];
p(6,:)= [2 2 4 5 4 5 5];
p(7,:)= [1 1 1 1 1 1 1]; %абсолютно худший продукт
p(8,:)= [5 5 5 5 5 5 5]; %абсолютно лучший продукт

%задаем значимость всех 7 критериев для фазы №3
%kr3=[z4 z3 z1 z2 z1 z2 z1]; %раздел 3
kr3=[z3 z4 z2 z3 z1 z2 z1]; %раздел 5

for i=1:8 %по каждому продукту вычисляем значение
    привлекательности
    pp=y(:,p(i,1));
    xpr0=fuzarith(x, pp, kr3(:,1), 'prod');
    t(i,1)= defuzz(x, xpr0, 'centroid');
    for j=2:7
        pp=y(:,p(i,j));
        xpr1=fuzarith(x, pp, kr3(:,j), 'prod');
        xpr2=fuzarith(x, xpr0, xpr1, 'sum');
        xpr0=xpr2;
        t(i,j)= defuzz(x, xpr1, 'centroid');
    end

    outp(i)=defuzz(x, xpr0, 'centroid');
    rp(:,i) = xpr0;
end
hold on;
```



```

plot(x, [rp(:,1)], 'r-', 'LineWidth',3);
plot(x, [rp(:,2)], 'r:', 'LineWidth',3);
plot(x, [rp(:,3)], 'r-.', 'LineWidth',3);
plot(x, [rp(:,4)], 'r--', 'LineWidth',3);
plot(x, [rp(:,5)], 'b:', 'LineWidth',2);
plot(x, [rp(:,6)], 'k--', 'LineWidth',2);
Legend ('C', 'MatLab', 'Deductor', 'Neural Designer', 'R',
'Python');
set(gcf, 'name', 'trapmf', 'numbertitle', 'off', 'Color','white')
ylim([0 1.2]);
display (prod);
display (outp);

```