

МОДЕЛЮВАННЯ ІННОВАЦІЙНИХ ІНТЕЛЕКТУАЛЬНИХ СИСТЕМ ПРИЙНЯТТЯ РІШЕНЬ В ЕКОНОМІЦІ

Ю. Г. Лисенко

Доктор економічних наук, професор, член-кореспондент НАНУ,
директор Навчально-наукового інституту
«Інноваційні технології управління»

Вищий навчальний заклад Укоопспілки
«Полтавський університет економіки і торгівлі»
вул. Коваля, 3, м. Полтава, 36014, Україна
yuriy.lysenko.1945@gmail.com

О. Ю. Мінц

Кандидат економічних наук, доцент,
докторант кафедри фінансів і банківської справи

Державний вищий навчальний заклад
«Приазовський державний технічний університет»
вул. Університетська, 7, м. Маріуполь, 87500, Україна
mints_a_y@pstu.edu

У статті систематизовано та розширено теоретико-методологічні засади процесів синтезу інноваційних інтелектуальних систем прийняття рішень. Визначено концептуальні підходи до застосування методів інтелектуальних обчислень для моделювання систем прийняття рішень. Запропоновано класифікацію методів інтелектуальних обчислень і класифікацію задач з аналізу та обробки даних. Розглянуто методологічні підходи до реалізації процесів спостереження, моделювання, ідентифікації та оцінки ефективності результатів моделювання інноваційних інтелектуальних систем прийняття рішень в економіці. Предметом дослідження є методологія моделювання інноваційних інтелектуальних систем прийняття рішень в економіці. Метою дослідження є формалізація процесів синтезу інноваційних інтелектуальних систем прийняття рішень для підвищення ефективності функціонування суб'єктів економічної діяльності. Результати дослідження дозволяють підвищити ефективність роботи із слабо структурованою інформацією, а також якість прийняття управлінських рішень в умовах невизначеності.

Ключові слова: *інтелектуальні обчислення, інтелектуальні системи, прийняття рішень, нейронні мережі, аналіз даних, обробка даних, ідентифікація, моделювання, оцінка ефективності.*

МОДЕЛИРОВАНИЕ ИННОВАЦИОННЫХ ИНТЕЛЛЕКТУАЛЬНЫХ СИСТЕМ ПРИНЯТИЯ РЕШЕНИЙ В ЭКОНОМИКЕ

Ю.Г. Лысенко

Доктор экономических наук, профессор, член-корреспондент НАНУ,
директор Учебно-научного института
«Инновационные технологии управления»

Высшее учебное заведение Укоопсоюза
«Полтавский университет экономики и торговли»
ул. Коваля, 3, г. Полтава, 36014, Украина
yuriy.lysenko.1945@gmail.com

А.Ю. Минц

Кандидат экономических наук, доцент,
докторант кафедры финансов и банковского дела
Государственное высшее учебное заведение «Приазовский
государственный технический университет»
ул. Университетская, 7, г. Мариуполь, 87500, Украина
mints_a_y@pstu.edu

В статье систематизированы и расширены теоретико-методологические основы процессов синтеза инновационных интеллектуальных систем принятия решений. Определены концептуальные подходы к применению методов интеллектуальных вычислений для моделирования систем принятия решений. Предложена классификация методов интеллектуальных вычислений и классификация задач анализа и обработки данных. Рассмотрены методологические подходы к реализации процессов наблюдения, моделирования, идентификации и оценки эффективности результатов моделирования инновационных интеллектуальных систем принятия решений в экономике.

Предметом исследования является методология моделирования инновационных интеллектуальных систем принятия решений в экономике.

Целью исследования является формализация процессов синтеза инновационных интеллектуальных систем принятия решений для повышения эффективности функционирования субъектов экономической деятельности.

Результаты исследования позволяют повысить эффективность обработки слабоструктурированной информации, а также качество принятия управленческих решений в условиях неопределенности.

Ключевые слова: интеллектуальные вычисления, интеллектуальные системы принятия решений, нейронные сети, анализ данных, обработка данных, идентификация, моделирование, оценка эффективности.

MODELING OF INNOVATIVE INTELLECTUAL DECISION-MAKING SYSTEMS IN THE ECONOMY

Yuriy Lysenko

Doctor of Economics, Professor, Corresponding Member of NASU,
Director of the Educational and Scientific Institute
"Innovative Management Technologies"

Higher Educational Establishment "Poltava University of Economics and Trade"
3 Kovalya str., Poltava, 36014, Ukraine
yuriy.lysenko.1945@gmail.com

Oleksii Mints

PhD in Economics, Docent,
Doctoral Candidate of the Department of Finance and Banking
State Higher Educational Establishment "Priazovsky State Technical University"
7 Universitetskaya str., Mariupol, 87500, Ukraine
mints_a_y@pstu.edu

The article systematizes and extends the theoretical and methodological foundations of the synthesis processes of innovative intellectual decision-making systems. Conceptual approaches to application of methods of intellectual calculations for modeling of systems of decision-making are defined. Classification of methods of soft computing, and classification of tasks of data analysis and data processing are offered. Methodological approaches to implementation of the processes of observation, modeling, identification and evaluation of effectiveness of the results of innovative intellectual decision-making systems modeling in the economy are considered.

The subject of study is the modeling of innovative intellectual decision-making systems in the economy.

The research objective is to formalize the synthesis processes of innovative intellectual decision-making systems to improve the efficiency of the functioning of economic systems.

The results of the research make it possible to improve the efficiency of poorly structured data processing and the quality of decision-making under uncertainty.

Keywords: *intellectual calculations, decision making, neural networks, data analysis, data processing, identification, modeling, assessment of efficiency.*

JEL Classification: C45, C51, C52, C81, D81

Вступ

Однією з глобальних сучасних тенденцій розвитку економіки є різке зростання кількості інформації, обсяги якої, за оцінками

експертів, підвищуються на 30 % щорічно [1]. 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. У зв'язку з цим 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 щодо використання таких технологій – «другий цифровий розрив». Тому проблема усунення цифрових розривів в Україні, де відставання в інноваційних технологіях обробки інформації рівнозначно відставанню в розвитку національної економіки, у даний час набуває особливої актуальності.

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. У моделі життєздатної системи С. Біра [2] застосування інтелектуальних систем прийняття рішень можливе у всіх підсистемах організаційної структури, що підвищує її адаптаційні здібності та життєздатність.

Різнорізані аспекти моделювання інтелектуальних систем прийняття рішень досліджуються в роботах вітчизняних і закордонних учених, серед яких варто згадати В. Вітлінського [3], В. Анфілатова [4], А. Матвійчука [5], С. Хайкіна [6], Г. Сетлак [7].

Разом із тим слід зазначити, що теоретичну базу та практичний інструментарій моделювання інноваційних інтелектуальних систем прийняття рішень розроблено недостатньо. У більшості

публікацій з даної тематики зосереджено увагу на використанні окремих методів, без намірів об'єднати їх у єдину систему. Тому комплексне дослідження питань моделювання інноваційних інтелектуальних систем прийняття рішень в економіці є актуальним науковим завданням.

Загальні визначення

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 підходи до класифікації та розробки методів інтелектуальних обчислень – «нисхідний» і «висхідний» [8].

Нисхідний, або семіотичний підхід (англ. *Top-Down AI*), має за мету 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*), полягає у створенні систем штучного інтелекту на основі моделювання базових біологічних і фізичних процесів. Результатом реалізації цього підходу є такі інструменти інтелектуальних обчислень, як штучні нейронні мережі.

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

Агентно-еволюційний підхід полягає у використанні для розв'язання 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 підхо-

ду можна виділити генетичні алгоритми та інші непереборні методи вирішення NP-повних задач (імітація відпалювання, метод мурашиних колоній тощо).

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

Класифікацію методів інтелектуальних обчислень за розглянутими підходами показано на рис. 1.



Рис. 1. Класифікація методів інтелектуальних обчислень

Як видно з класифікації, наведеної на рис. 1, деякі методи інтелектуальних обчислень можуть бути розглянуті в рамках кількох підходів. Але це не є наслідком недосконалої класифікації, а показує багатогранність інтелектуальних обчислень.

Концептуальні підходи до моделювання інноваційної інтелектуальної системи прийняття рішень

У загальному вигляді задача вибору рішення формулюється так [4]:

$$\omega = C(\Omega), \quad (1)$$

де Ω – множина можливих варіантів рішень; ω – обране рішення; C – правила вибору найкращої альтернативи (задаються у вигляді функції вибору).

Вибір відповідних методів пошуку рішень залежить від рівня визначеності множини варіантів Ω і ступеня формалізації функції вибору C , що можна бачити з табл. 1.

Таблиця 1

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
4.	Може доповнюватися	Не формалізовано	Загальна задача прийняття рішень	Нечітка логіка; Нейронні мережі; Методи логіко-лінгвістичного моделювання

З табл. 1 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 застосування таких підходів.

Якщо в задачі існує формальний критерій оптимальності, але характеристики простору рішень виключають можливість 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 штучні нейронні мережі Хопфілда, мережі з нейронами Поттса, зростаючі нейронні мережі, метод імітації відпалювання, метод мурашиних колоній та інші [9].

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 особи, що приймає рішення (ОПР). Для розв'язання задач вибору використовуються імітаційне моделювання, методи експертних оцінок, теорія корисності та інші підходи, які дозволяють зменшити невизначеність при виборі критеріїв.

Найхарактернішою задачею в управлінні складними системами є *загальна задача прийняття рішень* (ЗЗПР) [10]. Для неї характерна відсутність можливості визначення не тільки критеріїв оптимальності, але і те, що сам простір вибору постійно видозмінюється, що веде до необхідності безперервного моніторингу стану системи і вироблення коригувальних рішень. Підтримка прийняття рішень в ЗЗПР реалізується шляхом формування проміжної множини альтернатив, з яких здійснює вибір ОПР, тобто ЗЗПР зводиться до класу задач вибору. Формування проміжної множини альтернатив може здійснюватись за допомогою нейромережових інструментів, методів нечіткої логіки, а також методів логіко-лінгвістичного моделювання.

Для задачі пошуку оптимальних рішень, задачі вибору та загальної задачі прийняття рішень визначення гарантовано-кращого рішення за припустимий час неможливе внаслідок великої кількості можливих варіантів чи недостатньої формалізації правил вибору найкращої альтернативи. Такі економічні задачі будемо називати *складними*.

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 приводять в остаточному підсумку до вибору рішення, яке визначає подальший розвиток системи. У цій послідовності можна виділити процеси, пов'язані із спостереженням та моделюванням дослі-

джуваної системи, ідентифікацією її стану, оцінкою та вибором альтернатив.

Процес *спостереження* включає відбір із простору станів аналізованої системи Z множини характеристик системи Y , які спостерігаються, відстеження їх значень і збереження отриманої інформації в базі даних. У термінах системного аналізу рішення цієї задачі зводиться до відшукування такого відображення

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

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

Процес *моделювання* з формальної точки зору може розглядатися як побудова абстрактної множини E , ізоморфної предметної області:

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

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

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

Процес *оцінки та вибору альтернатив* для складних економічних задач пов'язаний із необхідністю згортки багатовимірного простору критеріїв [11]. При цьому кожній моделі поведінки економічної системи відповідає свій набір коефіцієнтів важливості критеріїв. Задача оцінювання альтернатив, таким чином, зводиться до відшукування відображення

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

тобто такого набору функцій μ_j , які б забезпечували однозначну відповідність між альтернативою $\omega \in \Omega$ та її оцінкою M відповідно до заданого критерію j .

Зазначені процеси можна звести у схему концепції моделювання інноваційної інтелектуальної системи прийняття рішень (ІСПР), представлену на рис. 2.

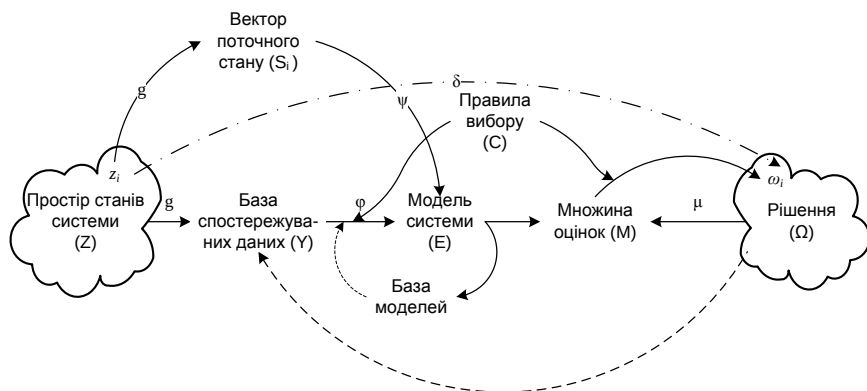


Рис. 2. Концепція моделювання інноваційної інтелектуальної системи прийняття рішень

Розглянемо процедуру формування рішення в ПСПР, схематично зображену на рис. 2.

Поточна ситуація z_i спостерігається та відображається у вигляді вектора S_i . Його структура повторює структуру елемента множини Y у частині вхідної інформації. Далі в процесі функціонування ПСПР 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, у розглянутій концепції ПСПР процес пошуку рішення зводиться до послідовності процесів спостереження, моделювання, ідентифікації, оцінювання та вибору. Розглянемо деякі аспекти, пов'язані з реалізацією цих процесів у ПСПР, які потребують ретельного вивчення та удосконалення.

Реалізація процесу спостереження

Будь-яка складна економічна система характеризується великою кількістю параметрів. У той же час кількість станів системи

(прикладів), які можуть бути відстежені за даними передісторії, звичайно обмежена. Тому в рамках процесу спостереження, як одного з етапів моделювання інноваційної інтелектуальної системи прийняття рішень, одними з головних є задачі, пов'язані з визначенням характеристик вхідної вибірки даних, зокрема її розмірності (кількості параметрів), необхідного обсягу (кількості прикладів) та значущості параметрів.

Проілюструємо зв'язок між цими характеристиками та ефективністю ідентифікації стану системи за допомогою графіка, наведеного на рис. 3 (дані умовні). Критерієм ефективності ідентифікації стану системи покладемо частку правильно розпізнаних прикладів на тестовій вибірці даних.

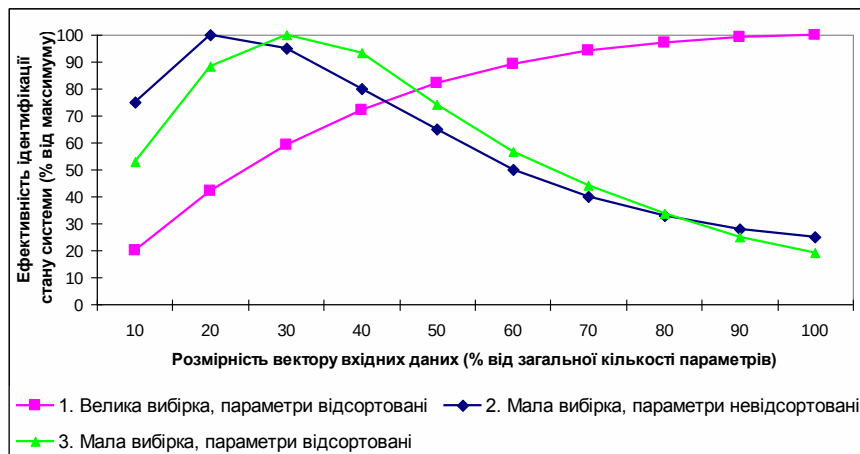


Рис. 3. Ілюстрація задачі пошуку оптимальних характеристик вхідних даних

Крива 1 на рис. 3 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.

Таким чином, на етапі спостереження перед дослідником можуть виникати різні проблеми, залежно від розмірності та обсягу вхідних даних. Так при вибірці великого обсягу, але низької розмірності, виникає проблема підвищення різноманітності даних, а при вибірці великої розмірності, але малого обсягу – проблема оцінки значущості параметрів. Тісно пов'язані з ними також проблеми визначення достатнього обсягу даних і синтезу оптимальної структури моделі, які повинні вирішуватися комплексно.

Хоча зазначені проблеми різною мірою стосуються багатьох методів інтелектуальних обчислень, найбільшу актуальність вони мають при використанні для аналізу даних штучних нейронних мереж (ШНМ), тому будуть розглянуті саме на їх прикладі.

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

$$T = \{X_i, Y_i\}_{i=1}^I, \quad (2)$$

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

$$X_i = (x_{i,1}, x_{i,2}, \dots, x_{i,M}), \quad (3)$$

де M – кількість параметрів, що складають один приклад.

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

$$NS = \langle N, L, \psi \rangle, \quad (4)$$

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. У цьому випадку структуру (точніше – архітектуру) мережі можна записати формулою виду

$$N_1 - N_2 - \dots - N_{lay}, \quad (5)$$

де lay – кількість шарів нейронної мережі; N_1 – кількість нейронів у вхідному шарі, що збігається з M – кількістю параметрів у навчальному прикладі; $N_2 \dots N_{lay-1}$ – кількість нейронів у кожному з прихованих шарів; N_{lay} – кількість нейронів у вихідному шарі, які є виходами ШНМ.

Так, наприклад, формулою 5-6-3-1 описується архітектура ШНМ, що має 5 входів, 6 нейронів у першому прихованому шарі, 3 нейрона в другому прихованому шарі та 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 забезпечувалась достатньо точна апроксимація функції $f(X_i)$ і, відповідно, рішення досліджуваної задачі. При цьому на вибірку $\bar{T}$ можуть накладатись обмеження залежно від умов задачі.

Існування проблеми відбору вхідних даних обумовлено залежністю між складністю структури нейронної мережі та кількістю прикладів, які необхідні для її навчання. Теоретично достатню кількість прикладів у навчальній вибірці можна визначити, відштовхуючись від концепції виміру Вапніка–Червоненкіса (VC-виміру) [12], що відображає «обчислювальну потужність сімейства функцій класифікації, реалізованих машинами, здатними до навчання» [6]. VC-вимір також може бути інтерпретований як «максимальне число образів, на яких штучна нейронна мережа може бути навчена без помилок для всіх можливих бінарних маркувань функцій класифікації» [6, с. 148]. Хоча точне аналітичне

визначення цього параметра для конкретної архітектури ШНМ у більшості випадків неможливе, оцінки, зроблені в [13], показують, що для найпоширенішого типу нейронної мережі (тобто мережі прямого поширення, яка складається з нейронів із сигмоїдальною активаційною функцією) верхня межа ВС-виміру пропорційна W^2 , де W – кількість вільних параметрів ШНМ, до яких відносяться вагові коефіцієнти зв'язків і параметри функцій активації нейронів:

$$W = |L| + |N|. \quad (6)$$

При цьому:

$$|L| = \sum_{i=1}^{lay-1} N_i \cdot N_{i+1}, \quad (7)$$

$$|N| = \sum_{i=2}^{lay} N_i. \quad (8)$$

Розглянемо вимоги до кількості прикладів у вхідній вибірці на прикладі порівняно невеликої повнозв'язаної нейронної мережі прямого поширення з сигмоїдальною активаційною функцією та архітектурою 5-5-1. У такій мережі $lay = 3$, отже з (6)–(8) $|L| = 25 + 5 = 30$, $|N| = 6 \Rightarrow W = 36$. Таким чином ВС вимір має той же порядок що $36^2 = 1296$. Це означає, що для найякіснішого налаштування параметрів даної нейронної мережі для ідентифікації стану складної системи достатньо вибірки, яка містить близько 1000 незалежних прикладів. Якщо кількість незалежних параметрів буде перевищувати ВС-вимір, це вже не поліпшить ефективність ідентифікації.

Інші методи оцінки (наприклад, метод, запропонований у [14]) дозволяють задавати допустимий рівень помилки мережі, що дає можливість отримати оптимістичніші значення розмірів навчальної вибірки при прийнятному рівні ефективності ідентифікації, однак вимоги до її обсягу все одно залишаються досить високими.

Слід наголосити на тому, що вираз W^2 обмежує лише теоретичну верхню межу ВС-виміру нейронної мережі [6]. Тому на

практиці може бути достатньо вибірки значно меншого розміру. Але вирази (6)–(8) дозволяють порівнювати різні варіанти архітектури нейронних мереж з погляду потреб до обсягу даних для навчання.

З (7) випливає, що в повнозв'язаних нейронних мережах кількість вільних параметрів 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 у навчальній вибірці навпаки – замала. 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 достатній обсяг вибірки на рівні 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 ефективного навчання нейронної мережі, не представляється можливим.

Проблема відбору вхідних даних тісно пов'язана з задачею знаходження такої структури мережі NS , яка б забезпечувала мінімум помилки моделювання. У сукупності ці дві задачі не мають і швидше за все не можуть мати аналітичних методів рішення. Не існує також і надійних емпіричних методів їх розв'язання, тому поширеним способом є простий перебір можливих варіантів. Однак, якщо при відносно невеликих обсягах вхідних 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 [15].

Розглянемо методи зниження розмірності даних, що ґрунтуються на:

- логічному аналізу даних;
- оптимізації представлення даних;
- кореляційному аналізу;
- непрямих методах аналізу значущості даних.

Очевидно, що основною метою застосування цих методів є видалення пояснюючих змінних, які слабо впливають або взагалі не впливають на результуючий показник. 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» 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. 4.

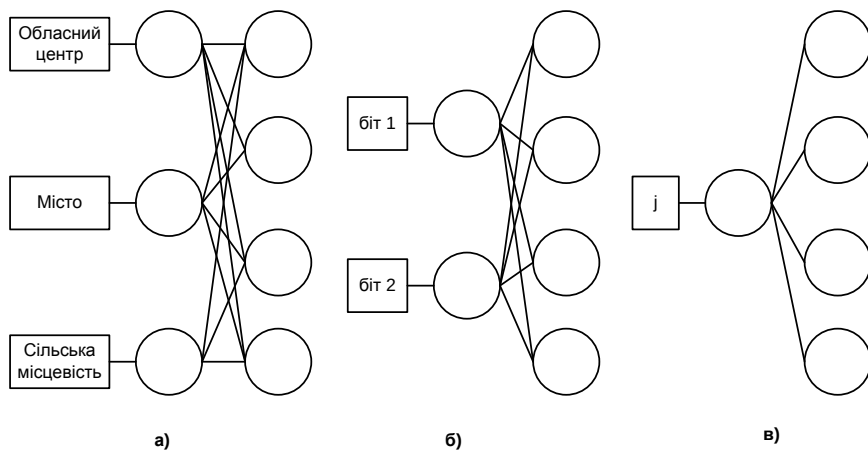


Рис. 4. Варіанти представлення нечислових даних, на прикладі параметра «Місце проживання клієнта»: а) представлення через позицію біта; б) представлення через бітову маску; в) представлення через ранг значення

Розглянемо переваги і недоліки кожного варіанта.

Варіант представлення даних через визначення *позиції біта* (рис. 4-а) для кожного значення параметра є найточнішим, але разом з тим і найбільш ресурсномістким (потребує найбільше параметрів ШНМ для його представлення).

Позначимо через v_j кількість можливих варіантів значення нечислового параметра j . При поданні параметра через позицію біта, відповідно до (6) – (8), буде потрібно

$$|J| = v_j \cdot N_2 \quad (9)$$

параметрів мережі для кодування цього показника, де N_2 – кількість нейронів у другому шарі ШНМ.

Для прикладу, який показано на рис. 4-а, $|J| = 3 \cdot 4 = 12$.

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. 4-б). Кількість параметрів мережі для подібного кодування нечислового показника, можна розрахувати за формулою

$$|J| = (\uparrow \log_2 v_j) \cdot N_2, \quad (10)$$

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

Для варіанту представлення даних з рис. 4-б вхідний показник може бути закодований двома бітами. При цьому значенню «Обласний центр» може відповідати маска «11», значенням «Місто» та «Село» – маски «10» та «01», відповідно. Кількість необхідних параметрів $|J| = 2 \cdot 4 = 8$.

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

Представлення нечислових параметрів у вигляді *рангів значень* (рис. 4-в) може використовуватись тоді, коли для цих параметрів існує критерій розподілу на шкалі «краще» – «гірше», але якщо такий критерій може бути синтезовано, виходячи з умов задачі. Якщо критерію розподілу немає, або він не є явним, даний спосіб кодування можна застосовувати тільки у виняткових випадках, що є основним недоліком даного способу. Перевагою його є мінімальні вимоги до ресурсів ШНМ. Дійсно, у цьому випадку

$$|J| = N_2, \quad (11)$$

що очевидно менше, ніж значення, отримані за формулами (9) або (10).

У розглянутому прикладі вхідні значення можна закодувати, наприклад, виходячи з питомої ваги проблемних кредитів в різних регіонах. Так, якщо найнадійнішими 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 та використовуваними ресурсами ШНМ. При великому обсязі вхідної вибірки або при малій кількості можливих значень параметра (наприклад, характеристика «стать» може приймати всього два значення) доцільно використовувати кодування позицією біта. За великої кількості можливих значень параметра або при невеликому обсязі навчальної вибірки доцільно використовувати кодування бітовою маскою. У окремих випадках може бути виправдане кодування шляхом ранжирування значень. Крім того, доцільно вивчити статистичні характеристики показників навчальної 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 результатом показником, так і параметри, які занадто сильно пов'язані з іншими вхідними факторами. 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.

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 сягає 1 тільки для лінійної залежності. Для кубічної та експоненційної залежності його

значення знаходиться в межах від 0,6 до 0,7, а для періодичних функцій, які зустрічаються в аналізі економічних систем досить часто, значення коефіцієнту кореляції близьке до 0 [16].

Значно кращих результатів із виявлення функціональних закономірностей дозволяє домогтися використання сучасних методів аналізу взаємозалежностей в даних. За результатами дослідження [16] найефективнішим виявився метод, заснований на обчисленні коефіцієнта максимуму взаємної інформації (maximal information coefficient, MIC). Значення MIC для всіх аналітично-заданих залежностей дорівнювало 1, що відповідає максимальному ступеню зв'язку.

Хоча коефіцієнт MIC було запропоновано порівняно недавно, у 2011 році він швидко набув поширення для аналізу слабоструктурованих даних. Вже у 2012 році К. Мерфі в монографії, присвяченій перспективам розвитку машинного навчання, назвав MIC кореляцією XXI століття [17, с. 61].

Слід зазначити і недоліки MIC. Головний з них обумовлений вимогами до дискретності аналізованих показників. Це змушує використовувати для аналізу неперервних величин алгоритми дискретизації даних, що негативно позначається на точності виявлення слабких залежностей. Для підвищення точності рекомендується вибір методу дискретизації проводити ітеративно, що дозволяє трохи підвищити точність але, в свою чергу, збільшує трудомісткість і витрати часу. Таким чином, використання MIC вимагає високої кваліфікації аналітика. Крім того, застосування цього методу стримується його відсутністю у поширених програмних продуктах.

Важливу роль при вирішенні задачі зниження розмірності вибірки можуть зіграти *непрямі методи аналізу значимості даних*. У роботі [18] 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 [19].

Використання цього алгоритму дозволяє провести ранжирування не тільки тих факторів, які мають числовий вираз, але і нечислових факторів, які не можуть бути приведені до числового

виду. Особливості роботи алгоритму дозволяють варіювати кількість параметрів, які відкидаються (тобто таких, яким присвоюється нульова значущість), що підвищує гнучкість методу. Ще однією перевагою алгоритму *S4.5*, порівняно з кореляційним аналізом, є автоматичне відкидання параметрів, що мають сильну взаємну кореляцію з іншими.

За малої розмірності вектора вхідних даних, великої кількості прикладів у вибірці та наявності складної залежності між вхідними і вихідними параметрами постає проблема *підвищення різноманітності вхідних даних*. Прикладом таких даних може бути біржова інформація з валютних торгів.

Базовий набір вхідних даних, що описує стан ринку, в загальному випадку містить лише інформацію про курсові відношення 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 у прогнозуванні навіть із застосуванням персептронних і подібних до них нейронних мереж, що є наслідком з фундаментальних обмежень, які були сформульовані Розенблаттом і доповнені М. Минським і С. Пейпертом [20]:

- 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 простіші;

- 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. Така інформація може бути отримана, наприклад, із використан-

ням емпіричних методів аналізу, які формуються у вигляді «умова» – «наслідок». Фактично, такі методи можна розглядати як мікро-експертні системи (ЕС-1 ... ЕС-п на рис. 5).

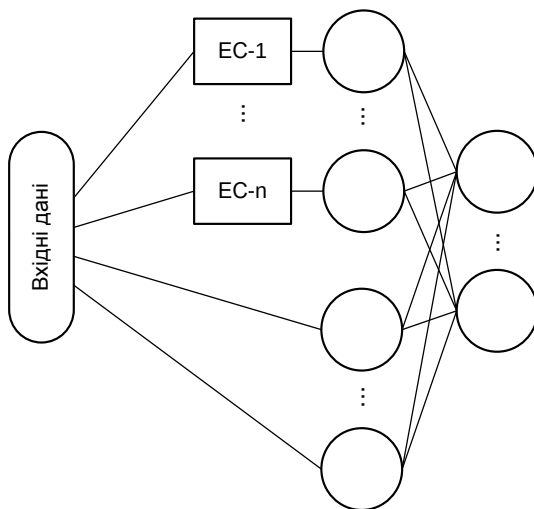


Рис. 5. Структура вхідної частини нейро-експертного модуля аналізу даних

Так, в якості ЕС-1...ЕС-п на рис. 5 на валютних і фондових ринках можуть виступати ринкові індикатори та осцилятори, які 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 [21].

Реалізація процесу моделювання

Серед методологічних проблем, які виникають на етапі моделювання ІСПР, особливо слід виділити такі, що пов'язані із постановкою завдання і вибором ефективних інструментів його рішення. Хоча в [5, 22] надаються деякі рекомендації щодо застосування певних нейромережових інструментів для вирішення різноманітних економічних задач, наприклад, побудови рей-

тингів, прогнозування, класифікації, проте не сформульовано єдиного комплексного підходу до конструювання штучних нейронних мереж з урахуванням специфіки задачі (вибору типу мережі, структури, обробки вхідних і вихідних даних тощо).

Аналіз використання цього терміна в різних літературних джерелах показує, що найбільше уваги визначенню правильної постановки задачі надається в рамках теорії розв'язання винахідницьких завдань, запропонованої та розвиненої Г. С. Альтшуллером і його послідовниками [23, 24]. В рамках цієї теорії постановка відокремлюється як самостійний етап розв'язання задачі, який значною мірою визначає остаточний результат.

За запропонованою концепцією моделювання інноваційних інтелектуальних систем прийняття рішень (див. рис. 2) постановка задачі є складовою частиною процесу моделювання ϕ . Результати проведеного дослідження показали, що при використанні інтелектуальних методів пошуку одна й та сама задача в деяких випадках може бути поставлена по-різному [25]. При цьому ефективність її рішення знаходиться у значній залежності від постановки задачі, а отже, від використаних методів та інструментів. Адже, відповідно даному вище визначенню складної економічної задачі, для неї неможливе знаходження гарантовано-кращого рішення за припустимий час. Тому всі методи розв'язання таких задач відшуковують тільки приблизні варіанти рішень. Тому ефективність застосування різних інструментів у загальному випадку відрізнятиметься.

Етап постановки задачі передбачає її віднесення до одного чи декількох класів. Тому слід розглянути класифікацію задач з аналізу та обробки даних.

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

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

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

Щоб уникнути можливих різночитань, далі термін «обробка даних» буде використовуватися виключно у вузькому сенсі.

Класифікацію задач з аналізу даних наведено на рис. 6.

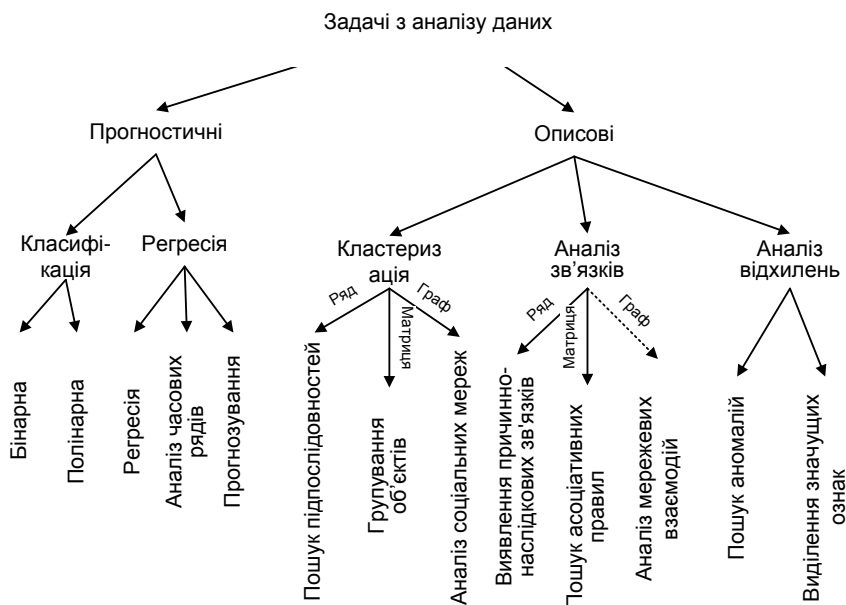


Рис. 6. Класифікація задач інтелектуального аналізу даних

На відміну від запропонованих раніше [26–30], класифікація на рис. 6 містить чотири ранги за рахунок виділення і впорядкування таксономічних ознак. Це дозволяє краще зрозуміти взаємозв'язок між різними класами задач і методами їх вирішення, що надає можливість підвищення ефективності аналізу даних.

Задачі аналізу даних за метою вирішення поділяються на дві великі групи – *прогностичні* та *описові*.

В рамках прогностичної групи на рис. 6 виділено такі класи:

1. *Задачі класифікації*. За кількістю класів, на які поділяється вхідна вибірка, слід виділяти задачі *бінарної* та *полінарної* класифікації. При бінарній класифікації вхідна вибірка ділиться тільки на два класи (наприклад – видавати кредит, або ні; надійний чи ненадійний контрагент). При полінарній класифікації вхідна вибірка ділиться на три або більше класів. Прикладами таких задач є: розпізнавання образів, сегментація клієнтів (якщо класи задані задалегідь) тощо. Незважаючи на схожість постановок задач бі-

нарної та полінарної класифікації, методи та інструменти їх вирішення суттєво різняться.

2. Задачі *регресії* потребують виявлення взаємозв'язку між вхідними та вихідними змінними. Прикладом задачі регресії може бути визначення суми кредиту, який може бути виданий клієнту. Існують також такі різновиди задач регресії, як прогнозування та аналіз часових рядів. Метою *прогнозування* є наближена оцінка значень деяких показників у майбутньому на підставі відомих значень у минулому і сьогодні. 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.

У задачах описової групи на рис. 6 виділено такі класи:

1. Задачі *кластеризації*. При їх вирішенні потрібно знайти закономірності в масиві даних, виділити в ньому деяку кількість зон (кластерів) і розподілити по ним дані. До цього класу віднесено задачі з *пошуку підпоследовностей у рядах даних, групування об'єктів та аналізу соціальних мереж*, які відрізняються тільки представленням вхідних даних і розмірністю простору угруповання. Так, пошук підпоследовностей у динамічних рядах фактично являє собою задачу кластеризації в одновимірному просторі часу, де кожен елемент має тільки двох сусідів – попереднє значення ряду і наступне. В задачах угруповання масиву вхідних даних кожен елемент (крім крайніх і кутових) має фіксоване число сусідів. У двовимірному ортогональному просторі їх чотири, в тривимірному – шість, і так далі. Вхідні дані у 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, яка в реальних соціальних мережах може сягати кількох тисяч (а для окремих людей і більше).

2. Задачі з *аналізу зв'язків* вирішуються за необхідності встановлення зв'язків і відносин між змінними у великих базах даних. Предметом *виявлення причинно-наслідкових зв'язків* є пошук статистично значущих закономірностей у часовій послідовності даних, який дозволяє відповісти на питання: «З якою ймовірністю настання події А тягне за собою настання події Б?». *Пошук асоціативних правил* дозволяє знаходити закономірності між

пов'язаними подіями, тобто дає можливість відповісти на питання: «З якою ймовірністю пов'язані події А і Б?». При цьому послідовність настання подій значення не має. *Аналіз мережевих взаємодій* дозволяє вирішити задачу пошуку епіцентрів мережевої активності, тобто вузлів, що впливають на процеси, які відбуваються в мережі, або ініціюють такі процеси. Вхідні дані при цьому подаються у вигляді графа. У широкому сенсі ця задача зводиться до виявлення для кожного вузла мережі *керуючих і керованих* вузлів.

3. До задач *аналізу відхилень* відносяться задачі з пошуку аномалій і виділення значущих ознак. *Пошук аномалій* включає виявлення та ідентифікацію таких елементів даних, які не відповідають встановленим 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 *виділення значущих ознак* набула особливої актуальності у зв'язку з розвитком Internet і різким зростанням обсягів інформаційних ресурсів. Суть її зводиться до створення на підставі вхідної інформації деякої вибірки, яка б при заданих обмеженнях на обсяг найповніше представляла її суть. Спочатку така задача вирішувалася виключно для текстових документів (автоматичне реферування), проте у даний час сфера її застосування охоплює всі основні види представлення інформації, у тому числі зображення, звук та відео. Серед актуальних прикладів її застосування – тематичний пошук контенту, контекстний пошук, виявлення матеріалів, що порушують інтереси правовласників чи законодавчі обмеження, тощо.

Класифікацію задач інтелектуальної обробки даних наведено на рис. 7.

Як видно з аналізу рис. 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 задачах зміна порядку елементів вхідних даних не відбувається.

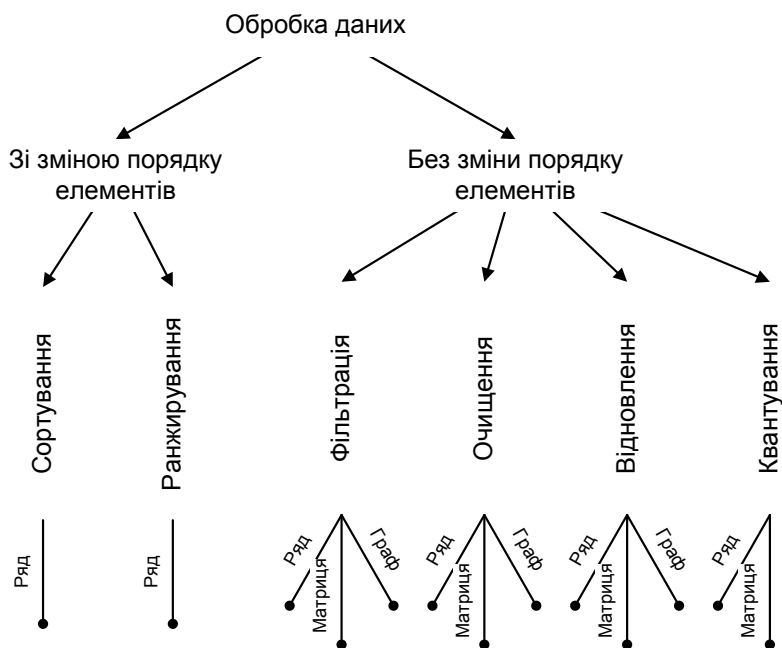


Рис. 7. Класифікація задач інтелектуальної обробки даних

Сортування – задача, пов’язана з розташуванням елементів даних у заданій послідовності або розподілом їх по групах.

Задача *ранжирування* відрізняється від сортування тим, що для її вирішення необхідно визначити метод визначення рангу кожного елемента вхідної послідовності даних (тобто метод, який дозволяє згорнути весь вектор значень елемента в єдиний параметр – ранг). Це дозволяє порівнювати між собою будь-яку кількість елементів із вхідного набору, не вдаючись до його повного сортування, що ефективно за наявності великих обсягів даних.

Під *фільтрацією* при обробці економічних даних будемо мати на увазі відбір інформації, що задовольняє заданому критерію. Вхідні дані для фільтрації можуть бути представлені у вигляді ряду, масиву або графа.

Задача очищення даних, по суті, є зворотною до задачі фільтрації та передбачає виключення з початкової вибірки «зайвих» даних. Очищення даних може стосуватись 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 у *матричній формі*, проводиться для усунення з вхідної вибірки факторів, що мало впливають на вихідні показники або зовсім не пов'язані з ними.

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

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 дозволяє підвищити ефективність машинного навчання за рахунок розширення вхідної вибірки даних. Необхідність відновлення даних може виникати при їх поданні у будь-якій 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.

Квантування використовується для динамічних рядів даних за необхідності зменшення кількості їх елементів, тобто при спрощенні рядів. Існують такі різновиди, 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 [31, с. 184]. У другому випадку розбиття динамі-

чного ряду на інтервали відбувається за часом. Кожен інтервал замінюється одним усередненим значенням елементів ряду, або декількома, що відображають граничні значення показника за цей період. Останній спосіб використовується для представлення біржових даних. Існують також різновиди квантування за рівнем і за часом зі змінним кроком квантування [32].

Складність задач обробки даних зростає з ускладненням структури поля критеріїв і зменшенням їх визначеності.

Таким чином, вирішення задач аналізу чи обробки даних у складі ПСПР доцільно розпочинати з постановки завдання та віднесення задачі до одного чи кількох класів із наступним моделюванням кожної з постановок та оцінкою ефективності результатів. Наприклад, задача біржового спекулянта може бути вирішена щонайменш у трьох постановках – класифікації, регресії та кластеризації, а задача передбачення банкрутств комерційних банків може бути зведена до постановок класифікації, прогнозування та кластеризації.

Реалізація процесу ідентифікації

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

Точність ідентифікації значною мірою визначається адекватністю моделі і, як наслідок, залежить від результатів процесу моделювання. Разом з тим, достовірність ідентифікації може бути

підвищена за рахунок зменшення невизначеності зовнішнього середовища шляхом оцінювання параметрів економічної системи, які на момент ідентифікації не є достовірно відомими. До них, зокрема, відносяться параметри, які формально визначаються за підсумками періоду, але можуть бути достатньо точно оцінені й раніше на підставі розрахункових чи імітаційних моделей. Останні розглянемо докладніше.

Існує кілька різновидів імітаційного моделювання [34, 35]:

– *Моделювання системної динаміки*. Виникнення цього напрямку пов'язане з ім'ям Дж. Форрестера, який у середині 1950-х років розробив його основні засади. Системна динаміка передбачає найвищий рівень агрегування компонентів з усіх методів імітаційного моделювання. Завдяки ряду спрощень, прийнятих в моделях системної динаміки (абстрагування від індивідуальних характеристик об'єктів і фізичних характеристик навколишнього середовища, неперервність усіх змінних і процесів), вони дозволяють простими засобами отримати адекватний опис процесів у досить складних системах. Модель системної динаміки може служити для виявлення та аналізу причинно-наслідкових зв'язків між будь-якими компонентами системи, дозволяє порівняти варіанти рішень з її управлінням. Наприклад, імітаційна модель ціноутворення на ринку житлової нерухомості [36] може бути використана для аналізу взаємозв'язків між факторами ринку і для прогнозування ринкових цін, що є важливим для кредитних організацій та ріелторів. Імітаційна модель сімейного бюджету банківського позичальника [37] дозволяє комерційним банкам оцінити ймовірність повернення кредитів і розрахувати можливі сценарії реструктуризації простроченої кредитної заборгованості.

– *Дискретно-подійне моделювання* передбачає дослідження функціонування системи в часі та аналіз впливу на її стан зовнішніх подій, які подаються у вигляді «заявок». Найвідомим 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 тощо.

– *Агентне моделювання*. Передбачає визначення поведінки одиничної простої структури – агента – у взаємодії з іншими такими ж агентами і навколишнім середовищем. Агент розглядається як деяка сутність, яка має активність, автономну поведінку, може приймати рішення відповідно до деякого набору правил, може взаємодіяти з оточенням та 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 та ін. Крім того, агентний підхід може застосовуватись і спільно з іншими різновидами імітаційного моделювання, зокрема в рамках моделей системної динаміки.

Таким чином, імітаційне моделювання відноситься до непрямих методів оцінки параметрів економічних систем і дозволяє розширити різноманітність вхідних даних для аналізу і, в результаті, підвищити обґрунтованість прийнятих рішень.

Реалізація процесу оцінки ефективності

Слід виділити як мінімум два підходи до оцінки ефективності інтелектуальних обчислень – абсолютний і відносний.

Визначення *абсолютних оцінок* дає можливість проаналізувати ефективність деякої моделі та відповісти на питання про її придатність або непридатність для вирішення поставленої задачі.

Визначення *відносних оцінок* дає можливість порівняння кількох моделей і визначення кращої з них за діючих умов.

Крім цього, задачу оцінки ефективності інтелектуальних систем прийняття рішень слід розглядати відносно різних об'єктів і процесів, серед яких відзначимо:

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

Розглянемо зміст поняття «ефективність» по відношенню до цих об'єктів.

Необхідність оцінки *ефективності роботи програмного забезпечення* обумовлена тим, що різні реалізації принципів інтелектуальних обчислень можуть сильно відрізнятися як за точністю, так і за швидкістю роботи. Аналогічна ситуація виникає при зіставленні ефективності роботи різних алгоритмів, що виконують функцію навчання моделі або налаштування параметрів інтелектуальних обчислень [38]. Таким чином, виникає потреба у надійному методі зіставлення різних програмних продуктів та алгоритмів за ефективністю їх реалізації.

Оцінку ефективності машинного навчання необхідно розглядати в зв'язку з вирішенням задачі оптимальної настройки алгоритмів навчання, які чутливі до обсягу навчальної вибірки та до параметрів процесу навчання. Існуючі підходи до оптимізації цих параметрів носять емпіричний характер і потребують узагальнення.

Ефективність моделей аналізу даних означає здатність розробленої моделі виконувати поставлені практичні завдання та має за мету визначення і порівняння економічної вигоди від їх застосування. Слід враховувати, що, принаймні, для різних класів економічних задач підходи до оцінки ефективності можуть відрізнятися.

При оцінюванні *ефективності моделей обробки даних* необхідно знайти відповідь на питання про економічну вигоду, що отримується в результаті застосування різних методів обробки.

Таким чином, задача дослідження 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 спричинив появу великої кількості модифікацій методів і алгоритмів машинного навчання. Так, кількість типів ШНМ на даний час перевищує 20. При цьому, наприклад, у математичному пакеті Matlab тільки для одного з цих типів – персептрона – передбачено більше ніж 10 варіантів навчання. Як правило, жодна з цих модифікацій не дозволяє отримати поліпшення ефективності на всьому діапазоні значущих параметрів і для всіх різновидів задач,

але для окремих видів приріст ефективності може бути суттєвим (див., наприклад, [38]). Аналогічна ситуація спостерігається і стосовно програмного забезпечення, виробники якого часто використовують спрощені реалізації методів та алгоритмів, які не дозволяють отримати найефективніші рішення.

Очевидно, що здійснення повного перебору всіх доступних варіантів для кожної задачі істотно збільшує трудомісткість реалізації ПСПР. Тому на практиці до цього вдаються рідко, покладаючись на досвід та інтуїцію розробника. Однак такий підхід не завжди дозволяє одержати найкращий результат. Ефективнішою є така організація процесу вибору, за якої безліч варіантів зводиться до мінімуму за допомогою апіорного порівняння.

Питання порівняння ефективності алгоритмів та їх реалізацій вперше було піднято Д. Кнотом [39]. Незважаючи на загальновизнану важливість таких досліджень, слід зазначити, що для оцінки ефективності роботи програмного забезпечення та алгоритмів інтелектуальних обчислень запропоновані в них ідеї повинні бути доповнені та розвинені з урахуванням новітніх розробок.

Зокрема, основним критерієм ефективності алгоритмів у [39] приймається швидкість роботи. Однак, з огляду на те, що слабкоструктуровані задачі не завжди можуть сходитись до глобального оптимуму похибки моделювання, іншим важливим критерієм слід вважати точність одержуваних результатів. Оскільки критерії швидкості та точності є взаємно суперечливими, залежно від умов задачі один із них доцільно обмежувати. Тобто, можна шукати або найточніший алгоритм при заданих обмеженнях за часом роботи, або найшвидший 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*.

Типовою назвемо задачу, основні характеристики якої з погляду вхідних даних і результатів є досить близькими для деякого набору інших задач у схожій постановці.

До задач, які можна використовувати за типові з економіки, висунемо такі вимоги (В.1 – В.6):

В.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 за різних умов.

В.2. *Прозорість економічної інтерпретації результатів*. Необхідно для забезпечення можливості прямого зіставлення алгоритмів, які досліджуються, за основним критерієм, прийнятим в економіці – вигоді від використання.

В.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.4. *Складність рішення*. Типові задачі мають відноситися до задач пошуку оптимальних рішень, задач вибору чи загальних задач прийняття рішень (див. табл. 1).

В.5. *Можливість порівняння результатів*. Постановка задачі повинна забезпечувати можливість визначення якості результатів різних її рішень.

В.6. *Репрезентативність*. Чим більше різних економічних задач може бути зведено до тих самих постановок, що і типова, тим краще.

Слід зазначити, що *de facto* у 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.1 – В.6 показує, що ця задача повністю відповідає лише вимогам В.3 і В.5, а також частково – В.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 [40, 41]. Суть її зводиться до відшукування найкоротшого шляху обходу заданих міст із подальшим поверненням на початок маршруту. Починаючи з 1950-х років ведеться систематичний пошук аналітичних рішень цієї задачі, що дозволило отримати методи і алгоритми розрахунку найкоротших шляхів для досить великої кількості вузлів. Недоліком цих методів є досить вузька спеціалізація, проте їх застосування дозволяє виконати вимогу В.3 до типових завдань. Інші вимоги також виконуються: дані легко можуть бути згенеровані (В.1); найкоротший маршрут є зазвичай і найвигіднішим економічно (В.2); знаходження найкоротшого шляху є дійсно складною задачею (В.4); відношення знайденого шляху до найкоротшого дозволяє однозначно порівнювати ефективність різних алгоритмів (В.5); до комбінаторної оптимізації в економіці можна звести велику кількість практичних задач (В.6).

Ще однією задачею, яка має велике значення для порівняння алгоритмів реалізації інтелектуальних обчислень, є задача біржового спекулянта, яка також відповідає вимогам В.1 – В.6. Так, результати біржових торгів є відкритими і загальнодоступними, що забезпечує виконання вимоги В.1. Отримані результати прямо виражаються через прибуток, що забезпечує прозорість їх економічної інтерпретації (В.2). Абсолютно кращий варіант рішення відповідає точному прогнозу розвитку подій, який відомий з аналізу даних передісторії. Це забезпечує виконання вимог з перевірки результату (В.3). Досягнення абсолютно кращого результату на практиці неможливе, оскільки на розвиток подій впливає безліч факторів, які лише побічно проявляються у вхідних даних (В.4). Результати роботи різних алгоритмів можна зіставити за розміром отриманого прибутку (В.5). Що стосується репрезентативності результатів (В.6), то вище вже зазначалось, що дана задача може розглядатись у різних постановках (регресії, класифікації, кластеризації), до яких можна звести велику кількість різних економічних задач з аналізу слабкозв'язаних даних.

Використання системи типових задач у поєднанні з розробленими раніше методами забезпечення порівняльності результатів [39, 42, 43] дозволяє підвищити обґрунтованість вибору програм-

них засобів і алгоритмів навчання та знизити витрати на створення ПСПР. При цьому система типових задач може формуватися поступово, на підставі узагальнення накопиченого досвіду в створенні ПСПР.

Ефективність машинного навчання. Термін «навчання» стосовно систем штучного інтелекту (комп'ютерних програм) визначено в класичній монографії Т. Мітчелла таким чином [44, с. 2]:

Кажуть, що комп'ютерна програма навчається при вирішенні якоїсь задачі з класу T , якщо її продуктивність, згідно метрики 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.

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

У *широкому* сенсі під навчанням будемо розуміти настройку самих навчальних алгоритмів, а також визначення структури системи інтелектуальних обчислень (наприклад – кількість нейронів у прихованих шарах ШНМ, кількість розгалужень у 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 у цьому випадку ділиться на навчальну та тестову, а якість навчання визначається за такими критеріями [45, с. 96]:

- К.1. Мінімізація помилки у розпізнаванні навчальної множини;
- К.2. Мінімізація помилки у розпізнаванні тестової множини;
- К.3. Досягнення адекватної динаміки процесу навчання.

Пояснимо критерій К.3. Аналіз динаміки навчання нейронної мережі дозволяє виявити момент перенавчання ШНМ, що є осно-

вною небезпекою роботи з малими вибірками даних, та уникнути його. На рис. 8 зображено класичний вид графіка зміни середньої помилки моделювання ШНМ на навчальній і тестовій множинах даних.

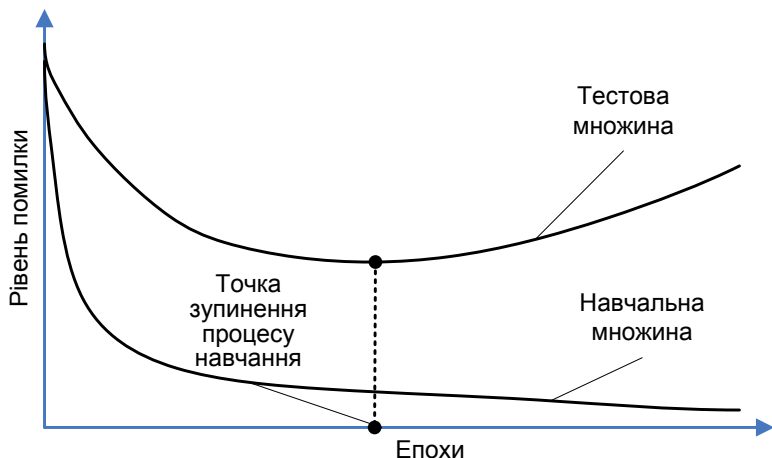


Рис. 8. Динаміка процесу навчання ШНМ

З аналізу рис. 8 видно, що якщо середня помилка на навчальній множині монотонно зменшується, то середня помилка ШНМ на тестовій множині проходить через екстремум, після чого починає збільшуватися. Проходження через екстремум і є моментом, коли процес навчання необхідно зупинити. У такому випадку ШНМ замість узагальнення залежностей між вхідними та вихідними даними починає підлаштовуватись і «запам'ятовувати» приклади з навчальної множини.

С. Хайкін вказує, що перехресна перевірка з використанням тестової множини і рання зупинка процесу навчання доцільні при виконанні такої умови [6, с. 293]:

$$I < 30 W, \quad (12)$$

де I – кількість прикладів у навчальній вибірці; W – кількість вільних параметрів нейронної мережі.

Тобто при малих обсягах навчальної вибірки перехресна перевірка необхідна.

Графік, показаний на рис. 8, є лише схематичною ілюстрацією процесу контролю за навчанням, оскільки за явної суперечливості характеристик навчальної вибірки і архітектури ШНМ його вигляд буде істотно відрізнятися від наведеного.

При аналізі ефективності навчання в задачах, де відсутня можливість розбиття вхідної вибірки на тестову і навчальну, застосування критерію К.2 неможливе. Якщо при цьому кількість ітерацій у процесі навчання не обмежена, може бути задіяний критерій К.3, відповідно до якого маркером закінчення процесу навчання слід вважати припинення поліпшення значень функції продуктивності. Наприклад, для генетичних алгоритмів такою буде функція пристосованості, а для самоорганізаційних ШНМ із шаром Кохонена – помилка кластеризації.

Відзначимо, що оцінка продуктивності інтелектуальної системи у вузькому сенсі за критерієм К.3 не є можливою для алгоритмів навчання з кінцевою кількістю ітерацій, наприклад для алгоритмів синтезу дерев рішень. Взагалі, непряма оцінка ефективності навчання за критеріями К.1 – К.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 оцінюванні результатів вирішення економічних задач головним критерієм ефективності є економічна вигода, тобто ефективність означає здатність розробленої моделі виконувати поставлені практичні завдання.

Оцінка ефективності отриманої моделі може здійснюватись 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. Однак перевірка ефективності на історичних даних дозволяє порівняти між собою значну кількість моделей і обрати для практичного використання ту, що має найбільшу очікувану ефективність, у той час як режим реальної експлуатації передбачає

аналіз зазвичай лише одної впровадженої системи. Тому обидва методи на практиці застосовуються спільно.

Розглянемо класифікацію моделей аналізу даних з погляду на оцінку їх ефективності, яка ґрунтується на класифікації, даних у [42, с. 565].

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.

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 відносяться, зокрема, ті, що вирішують задачі бінарної класифікації.

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

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. 8, можуть бути використані додаткові засоби візуальної оцінки ефективності прогнозування, зокрема *діаграми розсіювання*, що дозволяють візуально оцінити рівень і розподіл помилок прогнозування [42]. Вважається, що чим ближче розрахункові значення до реальних, тим менша помилка прогнозування. Однак така інтерпретація лише дублює показник середньоквадратичної помилки, тоді як *діаграми розсіювання* можуть використовуватись і для виявлення ефекту перенавчання.

Розглянемо *діаграми*, представлені на рис. 9. Вони отримані при аналізі розрахунків ШНМ, навчених на вибірці даних, об'єм якої згідно (6) і (12) можна трактувати як недостатній. Це викликає ефект перенавчання, в результаті чого нейронна мережа «запам'ятовує» вхідну вибірку даних і демонструє на ній дуже точні

результати прогнозування. На реальних же даних прогноз часто виявляється помилковим.

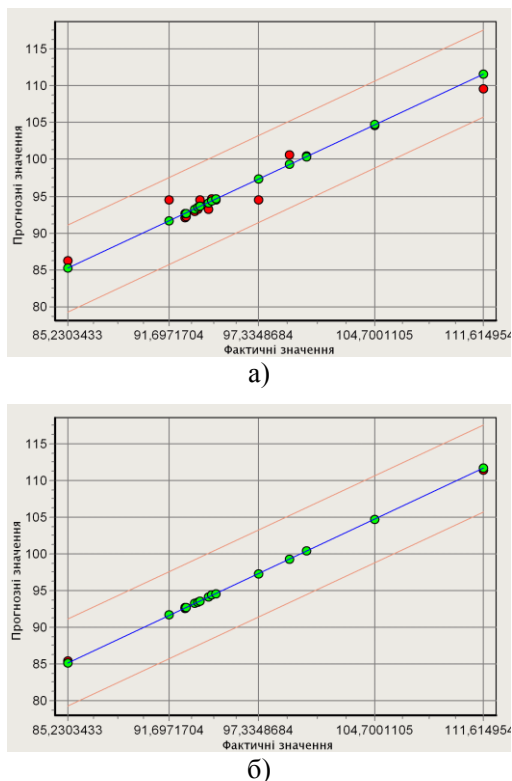


Рис. 9. Діаграми розсіювання моделей прогнозування
а) коректно навчена модель б) перенавчена модель

Для подолання ефекту перенавчання необхідно зменшити кількість вільних параметрів нейронної мережі до мінімуму, що забезпечує припустиму ефективність рішення задачі. Використання традиційного способу – розбиття вхідної вибірки на навчальну і тестову, з подальшим контролем помилки за тестовою вибіркою (див. рис. 8), у даному випадку може виявитись недостатньо ефективним з огляду на занадто малий розмір вхідної вибірки. У той же час, аналіз діаграм розсіювання різних моделей дозволяє отримати непряму інформацію щодо їх узагальнюючих здібностей.

Так, діаграма на рис. 9-а відображає результати, отримані за допомогою ШНМ, що містить у прихованому шарі 2 нейрона. Діаграма, показана на рис. 9-б, відображає результати розрахунків ШНМ, що містить у прихованому шарі 3 нейрона. Незважаючи на те, що діаграма на рис. 9-б ілюструє практично ідеальне розпізнавання вхідних прикладів, вона свідчить про перенавчання ШНМ. Тому для подальшого використання доцільно вибрати мережу з меншою кількістю нейронів, діаграма розсіювання якої наведена на рис. 9-а.

Ефективність кількісних моделей другого типу досліджується за допомогою інструментів, що дозволяють інтерпретувати результати інтелектуальних обчислень з погляду економічного ефекту та 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 тощо [42]. Розглянемо деякі з них.

Для формальної оцінки точності бінарного класифікатора розроблено різноманітні метрики оцінювання. Однак, найчастіше використовуються такі з них, як точність і повнота [43].

Точність класифікатора показує, скільки з передбачених позитивних результатів виявились дійсно позитивними:

$$Prec = \frac{TP}{TP + FP}, \quad (13)$$

де TP – кількість правильно розпізнаних позитивних результатів; FP – кількість негативних результатів, розпізнаних як позитивні.

Повнота класифікатора показує, скільки із загальної кількості позитивних результатів було передбачено правильно:

$$Rec = \frac{TP}{TP + FN}, \quad (14)$$

де FN – кількість позитивних результатів, розпізнаних як негативні.

Тут позитивним результатом є важливіший з точки зору задачі клас досліджуваних об'єктів. Скажімо, в задачі класифікації потенційних позичальників для банків важливіше точніше передбачати дефолти за кредитами, ніж визначати надійних позичальників, адже втрати за дефолтами критичніші, ніж недоотри-

мання прибутку від невидачі кредиту. Отже, у даному контексті дефолт стає позитивним результатом, у той час як виконання позичальником кредитних зобов'язань – негативним.

У різних практичних задачах одна з характеристик (*Prec* чи *Rec*) може виявитись важливішою за іншу. Наприклад, у задачах ранньої медичної діагностики важливіше повнота, а при класифікації потенційних банківських позичальників – точність.

Існують й інші формули для формальної оцінки бінарних класифікаторів, але часто наочнішим і зручнішим є візуальний аналіз. Одним з його інструментів є матриця спряженості, приклад якої наведено у табл. 2.

Таблиця 2

**ПРИКЛАД МАТРИЦІ СПРЯЖЕНОСТІ ПРИ АНАЛІЗІ
БАНКІВСЬКИХ ПОЗИЧАЛЬНИКІВ**

	Класифіковано		
Фактично	Позитивний	Негативний	Всього
Позитивний	34	1	35
Негативний	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 позитивних результатів як негативних, до *помилки другого роду* – класифікація негативних результатів як позитивних.

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.

Найскладнішими для оцінки ефективності результатів аналізу даних є *дескриптивні моделі*. Застосування апіорних методів оцінки економічного ефекту для них у більшості випадків немо-

жливе, 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, які описують дані з однаковою точністю, кращою є більш коротка [44]. 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, наведених вище на рис. 9, також відповідає ним.

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

$$e_{dp}^{k/h} = \frac{e_s^k}{e_s^h}, \quad (15)$$

де e_s^k та e_s^h – ефективність усієї ПСПР (або тої її частини, для якої вона може бути розрахована) за умови використання k -го та h -го методу обробки даних. При цьому інші методи та складові ПСПР залишаються незмінними (за винятком внутрішньої структури, створюваної в результаті навчання).

Оцінки, отримані за допомогою виразу (15), є відносними, тоді як на практиці зручнішими є абсолютні. Для отримання можливості роботи з абсолютними оцінками (щоб було зрозуміло, чи обробка даних дає позитивний ефект) можна використовувати концепцію «наївного» методу обробки, застосування якого гарантовано не покращує структуру вихідних даних. Для задач обробки даних *зі зміною порядку елементів* таким є випадковий вибір даних із 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 за формулою (15) дає результат більше 1, може вважатися ефективним, а

при результаті менше 1 – неефективним. При цьому виконується *правило транзитивності оцінок*, яке можна сформулювати так:

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

Існують також економічні задачі, вирішення яких збігається саме до задачі обробки даних (у цьому випадку процедура оцінки ефективності передбачає визначення економічного результату від застосування методу). Серед них є задача ранжирування, до якої зводиться широке коло економічних задач. Існують як суто розрахункові, так і розрахунково-графічні методи оцінки ефективності вирішення задач ранжирування. Серед останніх слід виділити Lift-криві та їх різновиди (Profit-криві, Gain-діаграми та ін.) [42].

Lift-крива формується на основі ліфт-фактора, який був вперше запропонований при вирішенні задачі оптимізації масової розсилки як показник, що відображає збільшення числа відгуків відносно кількості дій (поштових відправлень). За горизонтальною віссю графіка відкладається розмір вибірки, впорядкованої за зменшенням показника gp_i – імовірності настання позитивного результату, розрахованої аналізованою моделлю. За вертикаллю фіксується кумулятивне число позитивних результатів у кожній підвибірці (ліфт). Оскільки істинна рангова ознака trp_x у даному випадку є бінарною величиною, яка приймає значення 1 у разі позитивного результату x та значення 0 у разі негативного, вираз для розрахунку ліфта можна записати у такий спосіб:

$$lft(i) = \sum_{x=1}^i trp_x. \quad (16)$$

Як приклад використання Lift-кривої розглянемо задачу ранжирування прострочених кредитних справ у колекторському скорингу, в яких об'єкти необхідно розташувати за рівнем зменшення ймовірності відновлення позичальником платежів за кредитом. Припустимо, що для ранжирування застосовуються три моделі:

Модель 1 є емпіричною і використовується в багатьох кредитних підрозділах для ранжирування даних про позичальників. Ранговими ознаками gp_i в ній обрано кількість днів, які пройшли з моменту останнього платежу.

Модель 2 заснована на визначенні зв'язку між вхідними показниками та результирующим (імовірністю відновлення платежів) за допомогою логістичної регресійної моделі.

Модель 3 передбачає відсутність ранжирування, тобто випадковий вибір позичальника. Таку модель назвемо «*наївною*». Вона служить візуальним орієнтиром для аналізу та зіставлення ефективності інших моделей і завжди присутня на діаграмі з Lift-кривими.

Графіки Lift-кривих зазначених моделей наведено на рис. 10.

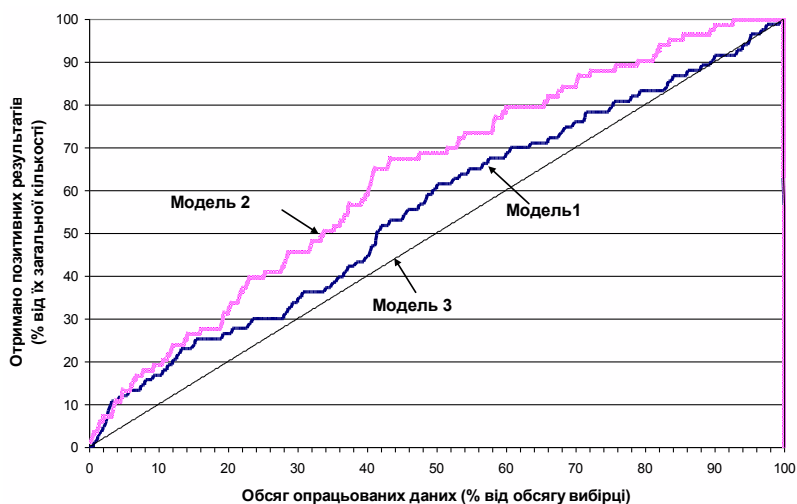


Рис. 10. Lift-криві різних моделей ранжирування позичальників у колекторському скорингу

Практичний сенс зіставлення ефективності за допомогою Lift-кривих полягає у визначенні моделі, що дозволяє зробити найменшу кількість дій, необхідних для досягнення певного результату.

Загальним критерієм оцінки моделі за її Lift-кривою є площа під кривою, що виражена у відсотках від загальної площини. Для «наївної» моделі цей показник завжди дорівнює 50 %. Для моделі логістичної регресії з рис.10 площа під кривою – 63,3 %. Для наведеної там же емпіричної моделі ранжирування – 55,1 %. Таким

чином, з розглянутих моделей за результатами зіставлення за площею під Lift-кривою ефективнішою для ранжирування позитивників виявилась модель логістичної регресії.

Іншим критерієм оцінки ефективності вирішення задач ранжирування є частка відгуків, одержуваних при здійсненні певної кількості дій. Так, з аналізу графіків на рис. 10 можна побачити, що для отримання 50 % позитивних відгуків у «наївної» моделі необхідно обробити 50 % кредитних справ, з використанням моделі логістичної регресії – 31 % кредитних справ, а з використанням моделі оцінки ймовірності погашення боргу за періодом прострочення – 41 % справ. Очевидно, що в даному випадку також доцільніше вибрати саме модель логістичної регресії. При цьому слід зазначити, що емпірична модель демонструє прийнятні результати при обробці малої кількості заявок. Отже, можливість дослідження не тільки загальної ефективності моделей, але 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 tr_x мають не бінарну природу. Однак, даний метод може бути вдосконалений для аналізу ефективності при довільній природі дійсних рангових ознак. Вираз для розрахунку ліфта при цьому:

$$lft(i) = \sum_{x=1}^i \{r \leq tr_x\}, \quad (17)$$

де $\{r \leq tr_x\} = 1$, якщо умова істинна, і 0 – якщо помилкова.

Метод побудови кривої ефективності ранжирування, заснований на функції (17), дозволяє забезпечити пріоритетну значимість правильного ранжирування об'єктів на початку списку. Дійсно, якщо за основний критерій ефективності моделі ранжирування прийняти площу під кривою (у частках від загальної площі графіка):

$$S_{lft} = \frac{\sum_{i=1}^n lft(i)}{n^2}, \quad (18)$$

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 у ранзі в запропонованому методі оцінки ефективності (17) убуває в арифметичній прогресії.

Висновки

Проведене дослідження дозволило запропонувати концепцію моделювання інноваційної інтелектуальної системи прийняття рішень, як сукупності процесів спостереження, моделювання, ідентифікації, оцінювання та вибору рішення. Розглянуто реалізацію цих процесів, що базується на розробках вітчизняних і зарубіжних авторів у таких областях, як статистичний аналіз, спостереження і експеримент, технології баз даних, математичне моделювання, інтелектуальних аналіз і обробка даних, методи оптимізації.

Вдосконалено та актуалізовано класифікацію економічних задач з аналізу й обробки даних та інтелектуальних методів їх вирішення.

Розвинуто методологію моделювання штучних нейронних мереж для вирішення економічних задач в частині вибору інструментів моделювання, визначення архітектури штучних нейронних мереж, вдосконалення методів роботи із вхідними даними, дослідження ефективності вирішення економічних задач методами нейромережевого моделювання.

Систематизовано та розширено підходи до оцінки результатів інтелектуальних обчислень, сформульовано та деталізовано концепцію типових задач, запропоновано розрахункові та розрахунково-графічні методи вирішення проблеми оцінки результатів обробки даних.

Розроблена методологія моделювання інноваційних інтелектуальних систем прийняття рішень в економіці надає можливість вирішити нагальну для управління економікою України пробле-

му подолання цифрових розривів, що забезпечуватиме підґрунтя для підвищення ефективності функціонування економічних систем різного рівня, відкриваючи нові можливості для економічного росту.

Література

1. Hilbert M. The World's Technological Capacity to Store, Communicate, and Compute Information / Martin Hilbert, Priscila López // Science. – 2011. – № 332 (6025). – Р. 60–65.
2. Бир С. Мозг фирмы / С. Бир. – М.: Едиториал УРСС, 2005. – 416 с.
3. Вітлінський В.В. Штучний інтелект у системі прийняття управлінських рішень / В.В. Вітлінський // Нейро-нечіткі технології моделювання в економіці. – 2012. – № 1. – С. 97–118.
4. Анфилатов В. С. Системный анализ в управлении / В. С. Анфилатов, А. А. Емельянов, А. А. Кукушкин. – М.: Финансы и статистика, 2002. – 368 с.
5. Матвійчук А. В. Штучний інтелект в економіці: нейронні мережі, нечітка логіка: монографія / А. В. Матвійчук. – К.: КНЕУ, 2011. – 439 с.
6. Хайкин С. Нейронные сети: Полный курс / С. Хайкин – [2-е изд.]. – М.: Вильямс, 2006. – 1104 с.
7. Сетлак Г. Интеллектуальные системы поддержки принятия решений / Г. Сетлак. – К.: ЛОГОС, 2004. – 250 с.
8. Copeland J. Artificial Intelligence: A Philosophical Introduction / Jack Copeland. – NJ: Wiley-Blackwell. – 1993. – 328 p.
9. Неітеративні, еволюційні та мультиагентні методи синтезу нечіткологічних і нейромережних моделей: Монографія / Під заг. ред. С. О. Субботіна. – Запоріжжя: ЗНТУ, 2009. – 375 с.
10. Сараев А. Д. Системный анализ и современные информационные технологии / А. Д. Сараев, О. А. Щербина. – Симферополь: СОНАТ, 2006. – 342 с.
11. Макаров И. М. Теория выбора и принятия решений / И. М. Макаров, Т. М. Виноградская, А. А. Рубчинский, В. Б. Соколов. – М.: Наука, 1982. – 328 с.
12. 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. – 1971. – № 17. – Р. 264–280.
13. Koiran P. Neural networks with quadratic VC dimension / P. Koiran, E. D. Sontag // Advances in Neural Information Processing Systems. – 1996. – № 8. – Р. 197–203.
14. Baum E. B. What size net gives valid generalization? / E. B. Baum, D. Haussler // Neural Computation. – 1989. – № 1. – Р. 151–160.

15. *Кравчук Е. В.* Искусственные нейронные сети и генетические алгоритмы / Е. В. Кравчук, Э. Хантер. – Донецк: ДонГУ, 2000. – 200 с.
16. *Reshef D. N.* Detecting Novel Associations in Large Data Sets / D. N. Reshef, Y. A. Reshef, H K. Finucane and others // Science. – 2011. – № 334 (6062). – P. 1518–1524.
17. *Murphy K. P.* Machine Learning: A Probabilistic Perspective (Adaptive Computation and Machine Learning series) – 1st Edition / Kevin P. Murphy. – Cambridge, MA: MIT Press, 2012. – 1067 p.
18. *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.
19. *Quinlan R. J.* C4.5: Programs for Machine Learning / Ross J. Quinlan // Machine Learning. – 1994. – Vol. 16. – № 3. – P. 235–240.
20. *Minsky M. L.* Perceptrons / M. L. Minsky, S. A. Papert. – Cambridge, MA: MIT Press, 1969. – 263 p.
21. *Эрлих А.* Технический анализ товарных и фондовых рынков / А. Эрлих. – М.: Юнити, 1996. – 318 с.
22. *Ежов А. А.* Нейрокомпьютинг и его применения в экономике и бизнесе / А. А. Ежов, С. А. Шумский., под ред. проф. В. В. Харитонов. – М.: МИФИ, 1998. – 224 с.
23. *Альтшуллер Г. С.* Творчество как точная наука / Г. С. Альтшуллер. – М.: Сов. радио, 1979. – 175 с.
24. *Альтшуллер Г. С.* Найти идею: Введение в ТРИЗ – теорию решения изобретательских задач. – 4-е изд. / Г. С. Альтшуллер. – М.: Альпина Пабlishерз, 2011. – 400 с.
25. *Минц А. Ю.* Общие вопросы постановки задач в нейросетевом моделировании / А. Ю. Минц // Нейро-нечіткі технології моделювання в економіці. – 2012. – № 1. – С. 189–206.
26. *Neelamadhab P.* The Survey of Data Mining Applications and Feature Scope / Padhy Neelamadhab, Mishra Pragnyaban, Rasmita Panigrahi // International Journal of Computer Science, Engineering and Information Technology. – 2012. – Vol. 2. – № 3. – P. 43–58.
27. *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. – 2013. – Vol. 13. – № 7. – P. 50–58.
28. *Sangameshwari B.* Survey on Data Mining Techniques In Business Intelligence / B. Sangameshwari, P. A. Uma // International Journal Of Engineering And Computer Science. – 2014. – Vol. 3. – № 10. – P. 8575–8582.
29. *Larose D. T.* Discovering Knowledge in Data: An Introduction to Data Mining / D. T. Larose. – NJ: Wiley & Sons, Inc, 2004. – 240 p.
30. Evolution of data mining, Gartner Group Advanced Technologies and Applications Research Note, 2/1/95 [Електронний ресурс]. – Режим доступу: <http://www.theartling.com/text/dmwhite/dmwhite.htm>.

31. Дьяконов В. Математические пакеты расширения Matlab. Специальный справочник / В. Дьяконов, В. Круглов. – СПб.: Питер, 2001. – 480 с.
32. Минц А. Ю. Метод упрощения динамических рядов с использованием генетических алгоритмов / А. Ю. Минц // Економічний вісник запорізької державної інженерної академії : зб. наук. праць. – Запоріжжє: ЗДІА, 2016. – Вип. 4 (04). – Ч. 2. – С. 120–124.
33. 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.
34. Карпов Ю. Имитационное моделирование систем. Введение в моделирование с AnyLogic 5 / Ю. Карпов. – СПб.: БХВ-Петербург, 2005. – 400 с.
35. Форрестер Дж. Основы кибернетики предприятия (Индустриальная динамика) / Дж. Форрестер. Пер. с англ. – М.: Прогресс, 1971. – 466 с.
36. Минц А. Ю. Моделирование ценообразования на рынке жилой недвижимости методами системной динамики / А. Ю. Минц // Технологический аудит и резервы производства. – 2016. – № 5/4(31). – С. 39–45.
37. Минц О. Ю. Моделювання процесів реструктуризації кредитів // Вісник Університету банківської справи НБУ: Зб. наук. праць. – К.: УБС НБУ, 2012. – № 2(14). – С. 329–333.
38. 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, 2007, March 7–9). – 2007. – P. 389–396.
39. Кнут Д. Искусство программирования. Том 1. Основные алгоритмы / Дональд Кнут. – М.: Вильямс, 2006. – 720 с.
40. Ватутин Э. И. Основы дискретной комбинаторной оптимизации / Э. И. Ватутин, В. С. Титов, С. Г. Емельянов. – М.: АРГАМАК-МЕДИА, 2016. – 270 с.
41. Карпенко А. П. Современные алгоритмы поисковой оптимизации. Алгоритмы, вдохновленные природой / А. П. Карпенко. – М.: Издательство МГТУ им. Н. Э. Баумана, 2014. – 446 с.
42. Паклин Н. Б. Бизнес-аналитика: от данных к знаниям / Н. Б. Паклин, В. И. Орешков. – СПб.: Питер, 2013. – 704 с.
43. 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. – 2011. – № 2 (1). – P. 37–63.
44. Mitchell T. Machine learning / T. Mitchell. – NY: McGraw-Hill, 1997. – 414 p.

45. Нечеткие модели и нейронные сети в анализе и управлении экономическими объектами: монография / под ред. Ю. Г. Лысенко. – Донецк: Юго-Восток, 2012. – 388 с.

References

1. Hilbert, M., & López P. (2011). The World's Technological Capacity to Store, Communicate, and Compute Information. *Science*, 332(6025), 60–65.
2. Bir, S. (2005). *Mozg firmy*. Moscow, Russia: Yeditorial URSS [in Russian].
3. Vitlinskiy, V.V. (2012). Shtuchnyi intelekt u systemi pryinyattya upravlynskykh rishen. *Neiro-nechitki tekhnologii modelyuvannya v ekonomitsi (Neuro-fuzzy modeling techniques in economics)*, 1, 97–118 [in Ukrainian].
4. Anfilatov, V. S., Yemel'yanov, A. A., & Kukushkin, A. A. (2002). *Sistemnyi analiz v upravlenii*. Moscow, Russia: Finansy i statistika [in Russian].
5. Matviychuk, A. V. (2011). *Shtuchnyi intelekt v ekonomitsi: neironni merezhi, nechitka logika*. Kyiv, Ukraine: KNEU [in Ukrainian].
6. Haykin, S. (2006). *Neyronnye seti: Polnyy kurs*. Moscow, Russia: Williams [in Russian].
7. Setlak, G. (2004). *Intellektualnye sistemy podderzhki prinyatiya resheniy*. Kyiv, Ukraine: LOGOS [in Ukrainian].
8. Copeland, J. (1993). *Artificial Intelligence: A Philosophical Introduction*. Oxford: Wiley-Blackwell.
9. Subbotin, S. O. (2009) *Neiteratyvni, evolyutsiyni ta multyagentni metody syntezy nechitkologichnykh i neyromerezhnykh modeley*. Zaporizhzhya, Ukraine: ZNTU [in Ukrainian].
10. Sarayev, A. D., & Shcherbina, O. A. (2006). *Sistemnyy analiz i sovremennyye informatsionnyye tekhnologii*. Simferopol, Ukraine: SONAT [in Russian].
11. Makarov, I. M., Vinogradskaya, T. M., Rubchinskiy, A. A., & Sokolov, V. B. (1982). *Teoriya vybora i prinyatiya resheniy*. Moscow, Russia: Nauka [in Russian].
12. Vapnik, V. N., & Chervonenkis, A. Ya. (1971). On the uniform convergence of relative frequencies of events to their probabilities. *Theoretical Probability and Its Applications*, 17, 264–280.
13. Koiran, P., & Sontag, E. D. (1996). Neural networks with quadratic VC dimension. *Advances in Neural Information Processing Systems*, 8, 197–203.
14. Baum, E. B., & Haussler, D. (1989). What sixe net gives valid generalization? *Neural Computation*, 1, 151–160.
15. Kravchuk, E. V., & Khanter, E. (2000). *Iskusstvennyye neyronnyye seti i geneticheskiye algoritmy*. Donetsk, Ukraine: DonGU [in Russian].
16. Reshef, D. N., Reshef, Y. A., Finucane, H. K., & others. (2011). Detecting Novel Associations in Large Data Sets. *Science*, 334(6062), 1518–1524.

17. Murphy, K. P. (2012). *Machine Learning: A Probabilistic Perspective*. Cambridge, MA: MIT Press.
18. Mints, A. Yu. (2011). Metody otbora dannykh dlya neirosetevogo modelirovaniya. *Modelyuvannya ta informatsiini systemy v ekonomitsi (Modelling and information systems in economy)*, 84, 256–270 [in Russian].
19. Quinlan, R. J. (1994, September). C4.5: Programs for Machine Learning. *Machine Learning*, 3(16), 235–240.
20. Minsky, M. L., & Papert, S. A. (1969). *Perceptrons*. Cambridge, MA: MIT Press.
21. Erlikh, A. (1996). *Tekhnicheskii analiz tovarnykh i fondovykh rynkov*. Moscow, Russia: Yuniti [in Russian].
22. Yezhov, A. A., & Shumskiy, S. A. (1998). *Neirokompyuting i yego primeneniya v ekonomike i biznese*. Moscow, Russia: MIFI [in Russian].
23. Altshuller, G. S. (1979). *Tvorchestvo kak tochnaya nauka*. Moscow, Russia: Sovetskoye radio [in Russian].
24. Altshuller, G. S. (2011). *Nayti ideyu: Vvedeniye v TRIZ – teoriyu resheniya izobretatel'skikh zadach*. Moscow, Russia: Alpina Publisherz [in Russian].
25. Mints, A. Yu. (2012). Obshchiye voprosy postanovki zadach v neirosetevom modelirovanii. *Neiro-nechitki tekhnologii modelyuvannya v ekonomitsi (Neuro-fuzzy modeling techniques in economics)*, 1, 189–206 [in Russian].
26. Neelamadhab, P., Pragnyaban, M., & Panigrahi, R. (2012). The Survey of Data Mining Applications and Feature Scope. *International Journal of Computer Science, Engineering and Information Technology*, 3(2), 43–58.
27. Seddawy, A. B. El, Turkey, S., & Ayman, K. (2013). Enhanced K-mean Algorithm to Improve Decision Support System under Uncertain Situations. *International Journal of Computer Science and Network Security*, 7(13), 50–58.
28. Sangameshwari, B., & Uma, P. A. (2014). Survey on Data Mining Techniques in Business Intelligence. *International Journal of Engineering and Computer Science*, 10(3), 8575–8582.
29. Larose, D. T. (2004). *Discovering Knowledge in Data: An Introduction to Data Mining*. Hoboken, New Jersey: Wiley & Sons.
30. Gartner Group Advanced Technologies and Applications Research Note. (1995). *Evolution of data mining*. Retrieved from <http://www.thearling.com/text/dmwhite/dmwhite.htm>.
31. D'yakov, V., & Kruglov, V. (2001). *Matematicheskiye pakety rasshireniya Matlab. Spetsial'nyy spravochnik*. St. Petersburg, Russia: Piter [in Russian].
32. Mints, A. Yu. (2016). Metod uproshcheniya dinamicheskikh ryadov s ispolzovaniyem geneticheskikh algoritmov. *Ekonomichnyi visnyk zaporizkoi*

derzhavnoi inzhenernoi akademii (Economic Bulletin of Zaporozhye State Engineering Academy), 4(2), 120–124 [in Russian].

33. Mints, A. Yu. (2015). Intellekturnyye metody analiza nadezhnosti uchastnikov rynkov finansovykh uslug. *Visnyk donetskogo universytetu ekonomiky ta prava (Bulletin of the Donetsk University of Economics and Law)*, 2, 85–90 [in Russian].

34. Karpov, Yu. (2005). *Imitatsionnoye modelirovaniye sistem. Vvedeniye v modelirovaniye s AnyLogic 5*. St. Petersburg, Russia: BKHV-Peterburg [in Russian].

35. Forrester, J. (1971). *Osnovy kibernetiki predpriyatiya (Industrial'naya dinamika)*. Moscow, Russia: Progress [In Russian].

36. Mints, A. Yu. (2016). Modelirovaniye tsenoobrazovaniya na rynke zhiloy nedvizhimosti metodami sistemnoy dinamiki. *Tekhnologicheskyy audit i rezervy proizvodstva (Technology audit and production reserves)*, 5/4(31), 39–45 [in Russian].

37. Mints, O. Yu. (2012). Modelyuvannya protsesiv restrukturyzatsii kredytiv *Visnik Universytetu bankivskoi spravy NBU (Bulletin of the University of Banking of the NBU)*, 2(14), 329–333 [in Ukrainian].

38. Schauerhuber, M., Zeileis, A., & Meyer, D. (2007, March 7–9). Benchmarking Open-Source Tree Learners in R/RWeka. *Proceedings of the 31st Annual Conference of the Gesellschaft für Klassifikation e.v. (Albert-Ludwigs-Universität: Freiburg)*, 389–396.

39. Knut, D. (2006). *Iskusstvo programmirovaniya. Tom 1. Osnovnyye algoritmy*. Moscow, Russia: Williams [In Russian].

40. Vatutin, E. I., Titov, V. S., & Yemel'yanov, S. G. (2016). *Osnovy diskretnoy kombinatornoy optimizatsii*. Moscow, Russia: Argamak-Media [in Russian].

41. Karpenko, A. P. (2014). *Sovremennyye algoritmy poiskovoy optimizatsii. Algoritmy, vdokhnovlennyye prirodoy*. Moscow, Russia: MGTU im. N. E. Bauman [in Russian].

42. Paklin, N. B., & Oreshkov, V. I. (2013). *Biznes-analitika: ot dannykh k znaniyam*. St. Petersburg, Russia: Piter [in Russian].

43. Powers, D. M. W. (2011). Evaluation: From Precision, Recall and F-Measure to ROC, Informedness, Markedness & Correlation. *Journal of Machine Learning Technologies*, 2(1), 37–63.

44. Mitchell, T. (1997) *Machine learning*. USA: McGraw-Hill.

45. Lysenko, Yu. G. (2012). *Nechetkiye modeli i neyronnyye seti v analize i upravlenii ekonomicheskimi ob'yektami*. Donetsk, Ukraine: Yugo-Vostok [in Russian].

Стаття надійшла до редакції 31.05.2017