

ПОБУДОВА АНСАМБЛІВ МОДЕЛЕЙ КРЕДИТНОГО СКОРИНГУ

Г. І. Великоіваненко

Кандидат фізико-математичних наук, професор,
завідувач кафедри економіко-математичного моделювання
Державний вищий навчальний заклад «Київський національний
економічний університет імені Вадима Гетьмана»
проспект Перемоги, 54/1, м. Київ, 03680, Україна
ivanenko@kneu.edu.ua

С. С. Савіна

Кандидат економічних наук, доцент,
доцент кафедри економіко-математичного моделювання
Державний вищий навчальний заклад «Київський національний
економічний університет імені Вадима Гетьмана»
проспект Перемоги, 54/1, м. Київ, 03680, Україна
s.savina@kneu.edu.ua

Д. В. Колечко

Член Правління, головний ризик-менеджер
Vietnam Prosperity Joint-Stock Commercial Bank (VPBank)
вул. Ланг Ха, 89, район Донг Да, м. Ханой, 117045, В'єтнам
dkolechko@gmail.com

В. П. Бенъ

Провідний спеціаліст
Акціонерне товариство «МОТОР СІЧ»
проспект Моторобудівників, 15, м. Запоріжжя, 69068, Україна
vlad.ben.1985@gmail.com

Стаття присвячена вирішенню актуального завдання підвищення ефективності оцінювання кредитних ризиків позичальників-фізичних осіб шляхом пошуку оптимального поєднання результатів розрахунків окремих скорингових моделей. Наведено принципи формування ансамблів моделей і проаналізовано існуючі підходи побудови ансамблевих структур. У процесі експериментального дослідження застосовано одну з модифікацій алгоритму бустінгу та реалізовано авторський алгоритм побудови ансамблю моделей на основі спеціалізації експертів. У якості окремих моделей-експертів використовувались нейромережі радіально-базисної архітектури. В результаті проведення порівняльного аналізу ефективності використаних ансамблевих технологій визначено, що найадаптованішою для задачі оцінювання кредитного ризику є запропонований авторами алгоритм побудови ансамблю на основі спеціалізації експертів.

Ключові слова: кредитний ризик, скорингова модель, радіально-базисна нейронна мережа, ансамбль (комітет) моделей, бустінг.

ПОСТРОЕНИЕ АНСАМБЛЕЙ МОДЕЛЕЙ КРЕДИТНОГО СКОРИНГА

Г. И. Великоиваненко

Кандидат физико-математических наук, профессор,
заведующий кафедры экономико-математического моделирования
Государственное высшее учебное заведение «Киевский национальный
экономический университет имени Вадима Гетьмана»

проспект Победы, 54/1, г. Киев, 03680, Украина
ivanenko@kneu.edu.ua

С. С. Савина

Кандидат экономических наук, доцент,
доцент кафедры экономико-математического моделирования
Государственное высшее учебное заведение «Киевский национальный
экономический университет имени Вадима Гетьмана»

проспект Победы, 54/1, г. Киев, 03680, Украина
s.savina@kneu.edu.ua

Д. В. Колечко

Член Правления, главный риск-менеджер
Vietnam Prosperity Joint-Stock Commercial Bank (VPBank)
ул. Ланг Ха, 89, район Донг Да, г. Ханой, 117045, Вьетнам
dkolechko@gmail.com

В. П. Бенъ

Ведущий специалист
Акционерное общество «МОТОР СИЧ»
проспект Моторостроителей, 15, г. Запорожье, 69068, Украина
vlad.ben.1985@gmail.com

Статья посвящена решению актуальной задачи повышения эффективности оценки кредитных рисков заемщиков-физических лиц путем поиска оптимального сочетания результатов расчетов отдельных скоринговых моделей. Приведены принципы формирования ансамблей моделей и проанализированы существующие подходы построения ансамблевых структур. В процессе экспериментального исследования применено одну из модификаций алгоритма бустинга и реализовано авторский алгоритм построения ансамбля моделей на основе специализации экспертов. В качестве отдельных моделей-экспертов использовались нейросети радиально-базисной архитектуры. В результате проведения сравнительного анализа эффективности использованных ансамблевых технологий подтверждено, что наиболее адаптированным для задачи оценивания кредитного риска является предложенный авторами алгоритм построения ансамбля на основе специализации экспертов.

Ключевые слова: кредитный риск, скоринговая модель, радиально-базисная нейронная сеть, ансамбль (комитет) моделей, бустинг.

BUILDING THE ENSEMBLES OF CREDIT SCORING MODELS**Halyna Velykoivanenko**

PhD (Physics and Mathematical Sciences), Professor,
Head of Department of Economic and Mathematical Modeling
State Higher Educational Establishment «Kyiv National Economic University
named after Vadym Hetman»

54/1 Peremogy Avenue, Kyiv, 03680, Ukraine
ivanenko@kneu.edu.ua

Svitlana Savina

PhD (Economic Sciences), Docent,
Associate Professor of Department of Economic and Mathematical Modeling
State Higher Educational Establishment «Kyiv National Economic University
named after Vadym Hetman»

54/1 Peremogy Avenue, Kyiv, 03680, Ukraine
s.savina@kneu.edu.ua

Dmitriy Kolechko

Member of Management Board, Chief Risk Officer
Vietnam Prosperity Joint-Stock Commercial Bank (VPBank)
89 Lang Ha str., Dong Da dist., Hanoi, 117045, Vietnam
dkolechko@gmail.com

Vladyslav Ben'

Leading Specialist
MOTOR SICH JSC
15 Motorobudivelnkyiv Avenue, Zaporizhzhia, 69068, Ukraine
vlad.ben.1985@gmail.com

The article is devoted to solving the actual problem of increasing the efficiency of assessing the credit risks of individual borrowers by finding the optimal combination of the results of calculations of specific scoring models. The principles of the formation of an ensemble of models are given and the existing approaches to the construction of ensemble structures are analyzed. In the process of experimental research has been applied one of the modifications of the boosting algorithm and implemented the author's algorithm for constructing an ensemble of models based on the specialization of experts. The radial-basis function neural networks were used as specific expert models. As a result of a comparative analysis of the efficiency of the used ensemble technologies it was confirmed that the algorithm for constructing an ensemble based on the specialization of experts proposed by the authors is the most adapted for the task of assessing credit risk.

Keywords: *credit risk, scoring model, radial-basis function (RBF) network, ensemble (committee) of models, boosting.*

JEL Classification: C45, D81, G21

І. ВСТУП

У процесі провадження банківської діяльності генеруються значні обсяги даних. Так, банки зберігають по кожному клієнту анкетні дані, кредитну історію, історію спілкування, фото- та відеоматеріали тощо. Крім власної інформації банківські установи часто отримують дані від бюро кредитних історій, операторів мобільного зв'язку, проводять моніторинг інтернет-активності клієнтів. Накопичення настільки різноманітних даних неодмінно призводить до проблем, пов'язаних з *Big Data*, коли виникає потреба врахування значної кількості суттєво неоднорідної інформації. Зокрема, при вирішенні одної з найактуальніших задач банківської діяльності – розробки системи скорингової оцінки кредитоспроможності позичальників – кількість характеристичних показників часто може сягати сотень і тисяч.

Проте, класичні рекомендації до побудови математичних моделей збігаються у висновку щодо доцільності застосування обмеженої кількості вхідних факторів. Дж. Міллер обґрунтовує оптимальну кількість входів моделей 7 ± 2 [1]. Так, якщо будується економетрична модель (а стандартом при побудові скорингових моделей у банківській сфері є логістичні регресії), то перевищення кількості вхідних факторів понад десять майже напевно призведе до прояву негативного явища мультиколінеарності. І хоча для нелінійних моделей (наприклад, нейронних мереж, які будуть використані у даній статті) це явище не несе прямих загроз, зростання кількості пояснюючих показників зменшує вплив кожного з них окремо на значення результативної змінної. За приблизно однакових результатів моделювання (точності відтворення вихідної статистики, прогнозування на незалежній тестовій вибірці та стійкості результатів прогнозування на різних вибірках) необхідно для подальшого використання брати ту економіко-математичну модель, яка має простішу структуру.

Якщо ж виникає необхідність врахування значної кількості факторів, то доречно модель зробити ієрархічною і пояснюючі змінні розподілити за різними узагальнюючими групами показників або джерелами інформації. Узагальнення результатів розрахунків отриманих моделей доцільно здійснити за допомогою ансамблевих технологій – одного з напрямків машинного навчання, що є ефективним засобом дослідження *Big Data* [2]. За такого підходу кожна з моделей може бути налаштована на маси-

вах однорідних даних і їх поєднання забезпечить коректне врахування всієї необхідної інформації.

Аналогічні міркування справедливі і стосовно обсягів спостережень у початковій сукупності даних. Так, складно розраховувати на ефективне налаштування однієї моделі на навчальній вибірці, що складається із сотень тисяч записів. Така модель не зможе відслідковувати усі наявні закономірності, оскільки вибірка міститиме інформацію, яка напевно буде поєднувати протилежні тенденції. У результаті цього зростатиме похибка агрегування, адже розрахунок такої моделі буде надто усередненим. Для уникнення подібних пасток моделювання доцільно розбити початковий масив даних на підвибірки, які утворюватимуться з прикладів із однотипними тенденціями поведінки. Для кожної такої вибірки будуватиметься окрема модель, узагальнення розрахунків яких реалізується шляхом використання ансамблів. На доцільності створення ансамблів наголошує і С. Терехов [3], який теоретично обґрунтовує та експериментально демонструє вищу точність вирішення задач класифікації на великих обсягах даних при об'єднанні розрахунків кількох моделей, побудованих для окремих сегментів даних, порівняно із застосуванням одного глобального класифікатора.

Метою роботи є вирішення актуального завдання підвищення ефективності оцінювання кредитних ризиків позичальників-фізичних осіб за рахунок пошуку механізму оптимального поєднання результатів розрахунків окремих скорингових моделей.

Досягнення цієї мети зумовлює необхідність вирішення таких завдань дослідження:

- проведення критичного аналізу основних підходів до оцінювання кредитоспроможності позичальників банків, виявлення їх недоліків і переваг;
- збір та аналіз інформаційної бази для розробки математичних моделей оцінки кредитного ризику позичальників банку;
- побудова математичних моделей оцінки кредитного ризику позичальників,
- дослідження різних підходів до формування ансамблів моделей та обґрунтування вибору раціональної ансамблевої структури для оцінювання кредитоспроможності позичальників;
- здійснення порівняльного аналізу результатів класифікації на основі застосування різних варіантів ансамблевих структур моделей.

II. ОГЛЯД ЛІТЕРАТУРИ

Історично першим дослідженням, присвяченим розробці ансамблів моделей, вважається праця Р. Шапайра 1990 року [4]. У цій статті викладено ідею бустінгу – алгоритму, що дає змогу підвищити точність класифікації однієї моделі. Подальші розробки в цьому напрямі Й. Фрюнда та Р. Шапайра 1997 року [5] привели до винаходу більш ефективної реалізації даного алгоритму, який отримав назву AdaBoost.

Надалі дослідження з розвитку ансамблевих технологій проходились у напрямках пошуку можливостей підвищення ефективності ансамблю. Зокрема, досліджувалась доцільність використання окремих видів моделей, наприклад, різних типів нейрон-мереж. Розглядалися різноманітні методи узагальнення результатів розрахунків моделей з метою мінімізації похибки роботи ансамблю на навчальних і тестових даних. Один із найпопулярніших на сьогодні алгоритмів побудови ансамблю моделей, що застосовується для розв'язання задач класифікації, запропоновано Л. Брейманом [6]. Однак, у подібних дослідженнях розглядаються загальні випадки вирішення задач моделювання (безвідносно моделей кредитного скорингу), тому вони мають більш теоретичне спрямування з огляду на мету нашого дослідження.

Лише останнім часом ансамблі моделей почали використовуватись для проведення скорингових оцінок, через що кількість публікацій за даною тематикою досить обмежена. Зокрема, у праці І. Кузнєцова та В. Кіреєва [7] описано процедуру розробки ансамблю для розв'язання задачі поведінкового скорингу фізичних осіб, наведено основні аспекти, що впливають на підвищення точності ансамблевих структур, та розглянуто реалізацію одного з них, а саме – метод узагальнення результатів розрахунків окремих моделей комітету. Відкритими для подальших досліджень залишаються інші питання конструювання ансамблевих архітектур, зокрема формування навчальних вибірок для побудови усіх моделей ансамблю.

III. ОБҐРУНТУВАННЯ ДОЦІЛЬНОСТІ ФОРМУВАННЯ АНСАМБЛІВ МОДЕЛЕЙ

При застосуванні ансамблів моделей одночасно з основною задачею розв'язується кілька простіших задач. Це обумовлюється тим, що досліджуваний повний масив інформації поділяється на се-

гменти, для кожного з яких розробляється окрема модель, а результати розрахунків моделей об'єднуються визначеним способом.

Однак не очевидним є факт, що результат об'єднання кількох класифікаторів дасть кращу якість, ніж окрема модель. Для дослідження даного питання розглянемо описаний С. Тереховим спрощений випадок роботи комітету з трьох моделей (які далі називатимемо експертами), що розв'язують задачу бінарної класифікації [3]. Позначимо ймовірності коректної класифікації (точність класифікації) кожним із них через p_1 , p_2 , p_3 . Ймовірності є незалежними. Нехай результат роботи комітету буде визначатись простим голосуванням, тобто певний приклад буде віднесено до того класу, який встановлено трьома чи двома експертами. Ймовірність правильної класифікації комітетом позначимо P_K .

У результаті роботи такого комітету можливі такі варіанти: всі три експерти провели класифікацію коректно, один з трьох помилився, помилились два з трьох та помилились всі три експерти.

Правильне рішення буде прийняте комітетом у двох перших варіантах. Деталізуючи їх отримаємо чотири можливі ситуації для правильного рішення:

- усі три експерти не помилились – ймовірність ситуації $\theta_1 = p_1 \cdot p_2 \cdot p_3$,

- помилився перший експерт, два інші провели класифікацію коректно – ймовірність ситуації $\theta_2 = (1 - p_1) \cdot p_2 \cdot p_3$,

- помилився другий експерт, два інші провели класифікацію правильно – ймовірність ситуації $\theta_3 = p_1 \cdot (1 - p_2) \cdot p_3$,

- помилився третій експерт, два інші правильно провели класифікацію – ймовірність ситуації $\theta_4 = p_1 \cdot p_2 \cdot (1 - p_3)$.

Ймовірність правильної класифікації комітетом буде складатись із сум ймовірностей чотирьох описаних ситуацій:

$$P_K = \theta_1 + \theta_2 + \theta_3 + \theta_4 = p_1 \cdot p_2 \cdot p_3 + (1 - p_1) \cdot p_2 \cdot p_3 + \\ + p_1 \cdot (1 - p_2) \cdot p_3 + p_1 \cdot p_2 \cdot (1 - p_3). \quad (1)$$

Якщо зафіксувати на певному рівні одну з ймовірностей, наприклад $p_3 = \text{const} = C$, то функція (1) перетворюється у функцію двох змінних. Змінюючи значення константи, можна наочно

відобразити поведінку функції P_K , досліджуючи таким чином результат роботи комітету моделей за різних значень якості класифікації окремими експертами.

На рис. 1 зображено вигляд поверхні, що відповідає виду функції P_K для випадку, коли $p_3 = C = 0,3$, а $0 \leq p_1, p_2 \leq 0,3$. За таких умов всі три експерти характеризуються низькою точністю класифікації, яка є гіршою, ніж можна було б отримати в результаті випадкового вибору (наприклад, при підкиданні монети).

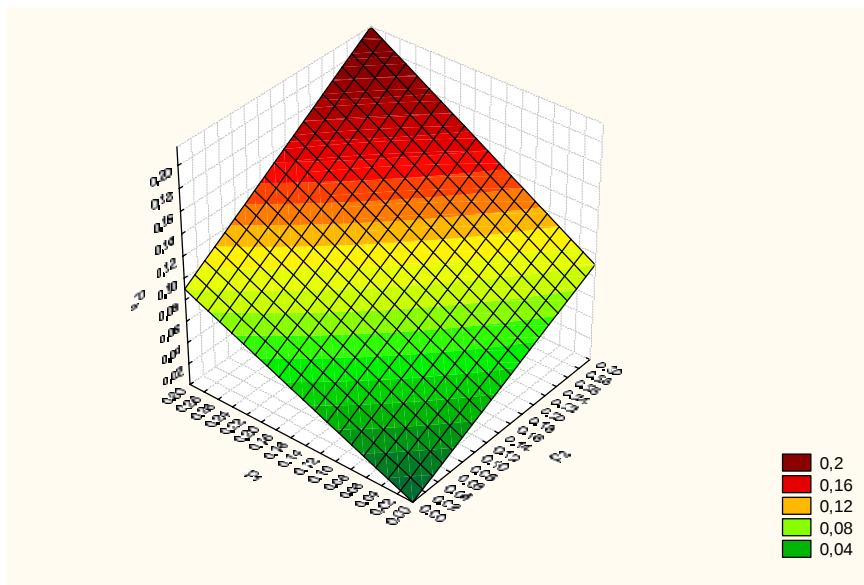


Рис. 1. Вигляд поверхні P_K при $0 \leq p_1, p_2 \leq 0,3; p_3 = 0,3$

На рис. 1 на поверхні виділено зони поступового зростання значення величини P_K , починаючи від 0 до максимального значення 0,216. Таке максимальне значення функції відповідає ймовірності правильної класифікації комітетом моделей, коли значення ймовірностей усіх трьох експертів окремо дорівнюють 0,3

(але, як бачимо, загальна точність є нижчою, ніж у кожного окремого експерта).

Дослідимо залежність імовірності правильної класифікації комітетом від точності класифікації двох моделей на повній множині значень імовірностей від 0 до 1 при фіксованій низькій точності третьої моделі на рівні 0,3. Так, у табл. 1 наведено значення, що приймає функція P_K у точках з дискретними координатами при зміні p_1 і p_2 від 0 до 1 з кроком 0,1 при $p_3 = 0,3$.

Таблиця 1

ЗНАЧЕННЯ ЙМОВІРНОСТІ КОРЕКТНОЇ КЛАСИФІКАЦІЇ КОМІТЕТОМ МОДЕЛЕЙ ДЛЯ ВИПАДКУ ФІКСОВАНОЇ ЙМОВІРНОСТІ КОРЕКТНОЇ КЛАСИФІКАЦІЇ ОДНОГО З ЕКСПЕРТІВ $p_3 = 0,3$

		Імовірність коректної класифікації другого експерта, p_2										
Імовірність коректної класифікації першого експерта, p_1		0	0,1	0,2	0,3	0,4	0,5	0,6	0,7	0,8	0,9	1,0
	0	0	0,03	0,06	0,09	0,12	0,15	0,18	0,21	0,24	0,27	0,30
	0,1	0,03	0,06	0,10	0,13	0,17	0,20	0,23	0,27	0,30	0,34	0,37
	0,2	0,06	0,10	0,14	0,17	0,21	0,25	0,29	0,33	0,36	0,40	0,44
	0,3	0,09	0,13	0,17	0,22	0,26	0,30	0,34	0,38	0,43	0,47	0,51
	0,4	0,12	0,17	0,21	0,26	0,30	0,35	0,40	0,44	0,49	0,53	0,58
	0,5	0,15	0,20	0,25	0,30	0,35	0,40	0,45	0,50	0,55	0,60	0,65
	0,6	0,18	0,23	0,29	0,34	0,40	0,45	0,50	0,56	0,61	0,67	0,72
	0,7	0,21	0,27	0,33	0,38	0,44	0,50	0,56	0,62	0,67	0,73	0,79
	0,8	0,24	0,30	0,36	0,43	0,49	0,55	0,61	0,67	0,74	0,80	0,86
	0,9	0,27	0,34	0,40	0,47	0,53	0,60	0,67	0,73	0,80	0,86	0,93
	1,0	0,30	0,37	0,44	0,51	0,58	0,65	0,72	0,79	0,86	0,93	1,00

За вказаних умов поверхня функції P_K є опуклою вгору, що можна бачити за лініями переходу кольорів на рис. 1, а також за розрахунками, наведеними в табл. 1. Найбільших значень поверхня набуває вздовж лінії $p_1 = p_2$, але при цьому значення функції $P_K < p_1 = p_2$. Отже, комітет у цьому випадку дає нижчу точність

оцінювання, ніж окремий експерт, що можна бачити за даними табл. 1. Така ситуація завжди характерна для роботи комітету, в якому хоча б одна з моделей має низьку точність класифікації. Перевищення точності оцінювання понад $P_K = 0,5$ такий комітет досягає лише за умови, що ймовірність коректної класифікації двома експертами з комітету буде не менше 0,6 (коли точність оцінювання третім експертом становить 0,3). А таку точність класифікації, яку має перший експерт, наприклад, на рівні 0,7, комітет моделей може продемонструвати лише при точності другого експерта 0,9.

На рис. 2 і 3 наведено випадки, коли ймовірність правильної класифікації третім експертом є вищою ($p_3 = 0,5$ і $0,7$, відповідно).

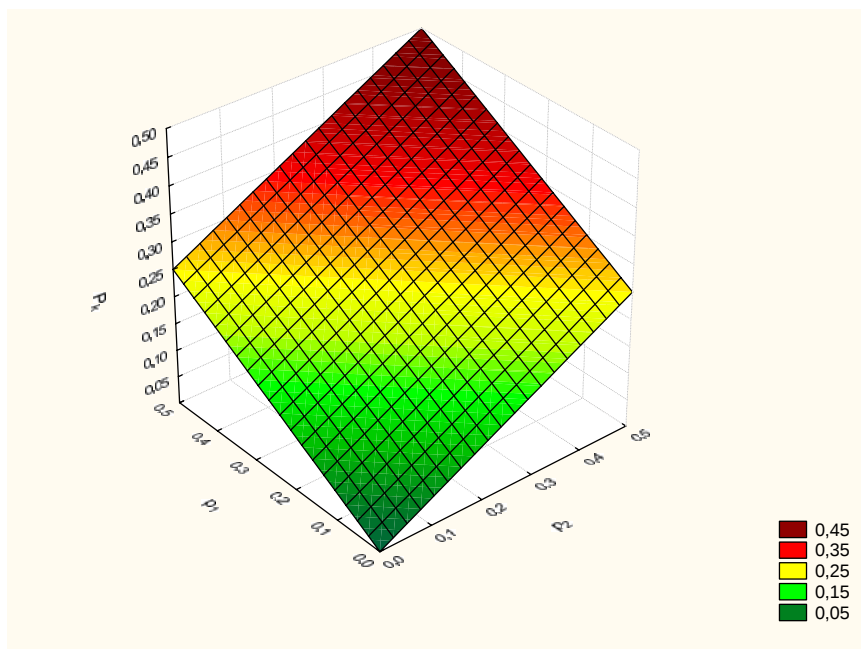


Рис. 2. Вигляд поверхні P_K при $0 \leq p_1, p_2 \leq 0,5; p_3 = 0,5$

На рис. 2 $p_3 = 0,5$ при $0 \leq p_1, p_2 \leq 0,5$. Вдovж лінії $p_1 = -p_2 + const$ при $p_3 = 0,5$ значення функції P_K є однаковими, як можемо бачити з даних, наведених у табл. 2. Крім того, рівні поверхні функції P_K зростають лінійно, що можна бачити за градієнтом переходу кольорів на рис. 2, тобто вона являє собою рівну площину. У даному випадку комітет лише не погіршує роботи окремих експертів.

Таблиця 2

**ЗНАЧЕННЯ ЙМОВІРНОСТІ КОРЕКТНОЇ КЛАСИФІКАЦІЇ КОМІТЕТОМ
МОДЕЛЕЙ ДЛЯ ВИПАДКУ ЙМОВІРНОСТІ КОРЕКТНОЇ КЛАСИФІКАЦІЇ
ТРЕТЬОГО ЕКСПЕРТА $p_3 = 0$**

Ймовірність коректної класифікації другого експерта, p_2												
Ймовірність коректної класифікації першого експерта, p_1		0	0,1	0,2	0,3	0,4	0,5	0,6	0,7	0,8	0,9	1,0
	0	0	0,05	0,10	0,15	0,20	0,25	0,30	0,35	0,40	0,45	0,50
	0,1	0,05	0,10	0,15	0,20	0,25	0,30	0,35	0,40	0,45	0,50	0,55
	0,2	0,10	0,15	0,20	0,25	0,30	0,35	0,40	0,45	0,50	0,55	0,60
	0,3	0,15	0,20	0,25	0,30	0,35	0,40	0,45	0,50	0,55	0,60	0,65
	0,4	0,20	0,25	0,30	0,35	0,40	0,45	0,50	0,55	0,60	0,65	0,70
	0,5	0,25	0,30	0,35	0,40	0,45	0,50	0,55	0,60	0,65	0,70	0,75
	0,6	0,30	0,35	0,40	0,45	0,50	0,55	0,60	0,65	0,70	0,75	0,80
	0,7	0,35	0,40	0,45	0,50	0,55	0,60	0,65	0,70	0,75	0,80	0,85
	0,8	0,40	0,45	0,50	0,55	0,60	0,65	0,70	0,75	0,80	0,85	0,90
	0,9	0,45	0,50	0,55	0,60	0,65	0,70	0,75	0,80	0,85	0,90	0,95
	1,0	0,50	0,55	0,60	0,65	0,70	0,75	0,80	0,85	0,90	0,95	1,00

Рис. 3, де $p_3 = 0,7$ при $0,5 \leq p_1, p_2 \leq 1$, відображає випадок, коли поверхня P_K є увігнутою. У даному випадку при $p_1 = p_2$

маємо $P_K > p_1 = p_2$. Отже, точність такого комітету дає результат кращий, ніж його окремі моделі.

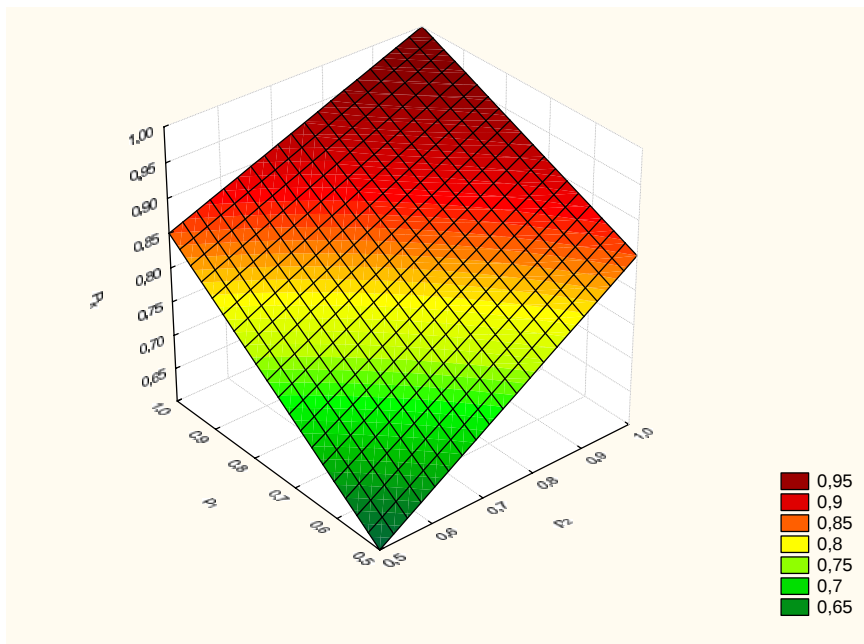


Рис. 3. Вигляд поверхні P_K при $0,5 \leq p_1, p_2 \leq 1; p_3 = 0,7$

Рис. 2 та 3 ілюструють головну умову доцільності створення комітету моделей: якщо точність хоч одного окремого експерта є нижчою за 0,5, то комітет за його участі стає менш ефективним у порівнянні з окремими моделями. Підвищення точності оцінювання комітетом моделей слід очікувати лише у випадку, коли ймовірність правильної класифікації є вищою за 0,5 для всіх експертів комітету [3].

Логіка представлених у тривимірному просторі поверхонь (рис. 1–3) може бути відображена в одному рисунку на площині (рис. 4).

Крива на рис. 4 відображає ймовірність точної класифікації всім комітетом. За горизонтальною віссю відкладається точність класифікації одним експертом (для спрощення приймається, що у всіх експертів однакова ймовірність p вгадування класу кредиту позичальників). Як бачимо на рис. 4, при значеннях ймовірності $p < 0,5$ точність комітету нижче бісектриси (тобто менше, ніж точність окремого експерта). У точці 0,5 крива лінія перетинає бісектрису, що вказує на рівність ймовірності точної класифікації комітетом і всіх трьох експертів. Далі крива проходить над бісектрисою – комітет демонструє вищу точність класифікації, ніж окремих експерт.

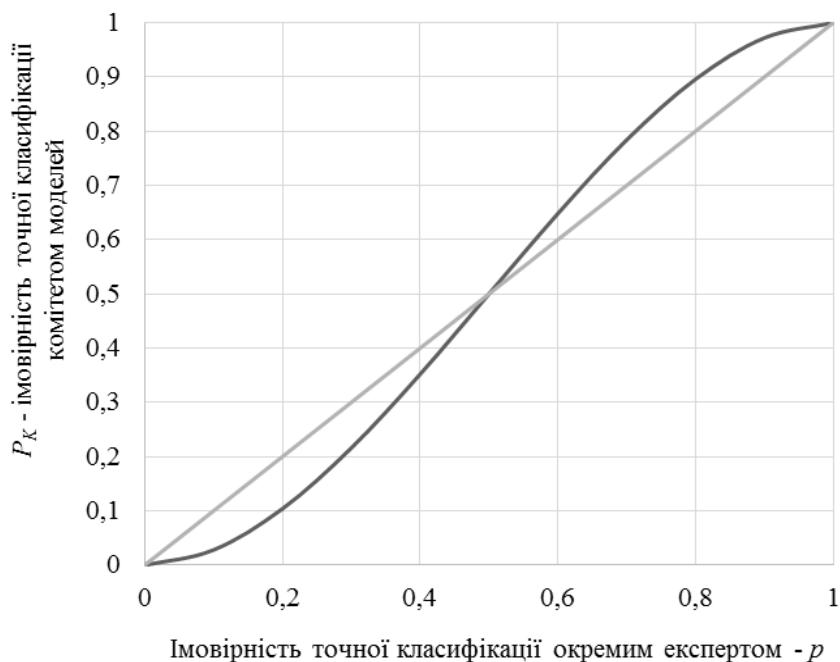


Рис. 4. Вигляд залежності ймовірності точної класифікації комітетом моделей P_K від точної класифікації одним експертом p (побудовано з використанням [3])

У праці [3, с. 15-16] зазначається, що логічним розвитком ідеї створення комітетів для підвищення точності класифікації є формування комітетів із комітетів, однак це не приводить до суттєвого зростання точності таких систем. Комітети лише адаптуються до інформаційного шуму в даних навчальної вибірки, однак точність на нових прикладах спадає (проявляється ефект перенавчання). С. Терехов у цій праці визначає дві умови зростання точності класифікації ансамблем моделей [3, с. 16]:

- необхідно підвищити точність класифікації кожної окремої моделі ансамблю,
- потрібно досягти статистичної незалежності похибок моделей ансамблю.

IV. ОСНОВНІ ВИДИ АНСАМБЛІВ

На сьогодні розроблено та описано значну кількість різноманітних видів ансамблів [3–10], які різняться за алгоритмами побудови. Деякі з них, наприклад бустінг, мають кілька модифікацій процесу побудови ансамблів і вже перетворились у окремі сімейства алгоритмів. Основні технології формування ансамблів і їх характеристики наведено в табл. 3.

Як бачимо за даними цієї таблиці, взагалі виділяють дві категорії ансамблів: зі статичною та динамічною інтеграцією моделей. При статичній інтеграції дані прикладу, для якого проводиться класифікація, не впливають на процедуру об'єднання результатів розрахунків моделей ансамблю. А при динамічній інтеграції процедура об'єднання результатів кожен раз коригується для проведення класифікації кожного конкретного прикладу. Крім того, алгоритми побудови ансамблів розрізняють також за способами формування навчальних вибірок для окремих моделей і схемою узагальнення результатів їх роботи.

Технології, розроблені для категорії статичних ансамблів, доцільно використовувати при дослідженні однорідних даних. До цих технологій відносяться усереднення, стекінг, бустінг, беггінг.

Таблиця 3

АЛГОРИТМИ ФОРМУВАННЯ АНСАМБЛІВ МОДЕЛЕЙ ТА ЇХ ОСНОВНІ ХАРАКТЕРИСТИКИ

Категорії ансамблів	Статична інтеграція моделей ансамблю			Динамічна інтеграція моделей ансамблю		
	Збереження початкової щільності розподілу даних	Стекінг	Бустінг	Беггінг	Суміш думок експертів	Ієрархічне поєднання думок експертів
Наявність впливу на початкову щільність розподілу даних						
Назва алгоритму формування ансамблю	Зміна початкової щільності розподілу даних	Зміна початкової щільності розподілу даних	Зміна початкової щільності розподілу даних	Зміна початкової щільності розподілу даних	Зміна початкової щільності розподілу даних	Зміна початкової щільності розподілу даних
Технологія формування початкової вибірки	Навчальна вибірка єдина для всіх моделей ансамблю, навчання проходить паралельно	Навчальна вибірка єдина для всіх моделей ансамблю, об'єкти змінюються в процесі навчання різних типів моделей	Для кожної наступної моделі ансамблю використовується нова навчальна вибірка, у якій змінюються розподіл даних, виходячи з результатів попередніх моделей	З однієї навчальної вибірки за допомогою спеціальної процедури випадкового вибору з поверненням ігнорування утворюються кілька навчальних вибірок для побудови моделей	Початкова вибірка поділяється на класи, які є навчальними вибірками для всіх моделей ансамблю, зберігаються також оцінки помилок експертів на кожній навчальній вибірці	Початкова вибірка поділяється на класи, які є навчальними вибірками для всіх моделей ансамблю, зберігаються також оцінки помилок експертів на кожній навчальній вибірці
Схема об'єднання результатів окремих моделей	Усереднення результатів по ансамблю	Застосування метамоделі	Усереднення результатів по ансамблю	Усереднення результатів або загальне усереднення по ансамблю	При класифікації нового об'єкта спочатку визначається найближчий до нього кластер, а потім використовуються ті моделі, які мали найменші помилки в даному кластері (найпростіший спосіб), або остання результатів моделей проводиться з вагами, які є виходом деякої керуючої метамоделі	Застосування керуючих метамоделей на двох рівнях ієрархії

Технологія усереднення результатів роботи моделей ансамблю є найпростішою. Вона передбачає, що всі моделі було побудовано на одній навчальній вибірці, а результат роботи комітету визначається простим голосуванням, тобто більшістю голосів.

Розташування алгоритму стекінгу в табл. 3 слід вважати дещо умовним, оскільки Н. Паклін і В. Орешков [8] зазначають, що загальна концепція використання даного методу відсутня, а його головна ідея знаходить свою реалізацію в різноманітних варіантах. Сутність даного алгоритму полягає в такому:

1) для комбінування в ансамбль застосовуються моделі не одного, а різних типів – наприклад, нейромережа, дерево рішень і логістична регресія;

2) замість алгоритму голосування вводиться концепція метанавчання.

На вхід деякої метамоделі, яку називають моделлю першого рівня, подаються результати класифікації кожним експертом, які відносять до нульового рівня. Далі результати розрахунків моделі першого рівня передаються на модель другого рівня і так далі, доки не буде досягнуто певної умови зупинки процесу.

Найбільш розвиненими та поширеними алгоритмами формування комітетів моделей є бустінг і беггінг.

За алгоритмом бустінгу моделі ансамблю будуються послідовно таким чином, щоб кожна наступна модель проводила класифікацію тих прикладів, які не були коректно ідентифікованими моделями на попередніх кроках.

На відміну від бустінгу, за алгоритмом беггінгу моделі будуються паралельно та незалежно одна від одної [8]. Так, на даних вихідного масиву шляхом випадкового відбору формується кілька вибірок. Вони матимуть той самий розмір, що і початковий масив, однак за рахунок випадкового відбору набір прикладів у кожній вибірці буде різним – один і той самий приклад може потрапляти в одну вибірку як кілька разів, так і жодного разу. Далі для кожної вибірки будується окрема модель-експерт. Результати розрахунків узагальнюються голосуванням.

Для масивів неоднорідних даних окремі їх підмножини можуть краще описуватись різними моделями, тобто побудову моделей-експертів слід здійснювати не для всієї вибірки в цілому, а для окремих її частин. У таких випадках мова йде про спеціалізацію експертів. Для врахування спеціалізації при формуванні ансамблів використовуються технології динамічної інтеграції моделей.

Алгоритм побудови ансамблю для неоднорідних даних пропонується у праці [3, с. 33]. За цим алгоритмом необхідно провести кластеризацію даних із застосуванням карт Кохонена і для кожного кластера обрати той набір моделей, який демонстрував для нього найвищу ефективність. Тоді для класифікації нового прикладу потрібно буде визначити найближчий кластер і застосувати відповідні моделі.

Авторами цієї статті було проведено дослідження [11, 12] з використанням ансамблю моделей для оцінки кредитоспроможності позичальників-фізичних осіб, який побудовано за алгоритмом усереднення результатів моделей. Для розрахунків використовувалась база даних з 2175 спостережень, яка містить широкий набір характеристичних показників позичальників комерційного банку та відомостей про виконання ними зобов'язань за отриманими кредитами.

З метою відбору найбільш значущих чинників для оцінювання кредитоспроможності позичальників було використано підхід, що заснований на поєднанні роботи ймовірнісної нейромережі та генетичного алгоритму. Також було застосовано два більш спрощених підходи на основі покрокового включення та покрокового виключення чинників з моделі. Узагальнення результатів моделювання за трьома підходами дозволило виділити з множини початкових пояснюючих чинників шість кількісних (вік, стаж на останньому місці роботи, загальний стаж, наявність депозитів, наявність виплачених у минулому кредитів і кількість дітей у сім'ї) та два якісних (рівень освіти та статус працюючого), які було обрано як вхідні змінні моделей оцінювання кредитоспроможності позичальників.

На основі цих факторів у праці [12] було сформовано ансамбль з трьох нейромереж: дві мережі з радіально-базисною архітектурою (82 та 124 нейрони на прихованому шарі) та тришаровий персептрон з 6 нейронами проміжного шару. Всі моделі мали єдину навчальну вибірку, їх навчання проходило незалежно, загальний розрахунок здійснювався за принципом простого голосування. Межі зміни точності класифікації отриманим ансамблем і його окремими моделями в залежності від умов навчання наведено в табл. 4.

Як бачимо з даних, наведених у табл. 4, точність класифікації комітетом моделей не перевищує точності найефективніших за даних умов моделей. Однак зміна умов експерименту приводить до більших коливань в ефективності окремих моделей, ніж сформованого з них комітету.

Таблиця 4

**МЕЖІ ЗМІНИ ТОЧНОСТІ КЛАСИФІКАЦІЇ АНСАМБЛЕМ
І ЙОГО ОКРЕМИМИ МОДЕЛЯМИ ЗАЛЕЖНО ВІД УМОВ НАВЧАННЯ, %**

Вид класифікатора	Відсоток правильно класифікованих спостережень	
	у навчальній вибірці	у тестовій вибірці
Ансамбль моделей	60,7–68,5	47,5–52,3
Окремі нейромережі	55,4–71,9	43,7–55,2

Зазначимо, що невисокі показники точності класифікації обумовлюються специфікою застосованого у дослідженні пакету *STATISTICA*, в якому границя поділу між класами встановлена на рівні 0,5 (вищі значення розрахунку моделі інтерпретуються як такі, що належать до класу «1», а нижчі значення – до «0»). Однак, проведені нами численні експериментальні розрахунки із різними типами моделей свідчать, що границя раціонального поділу між класами (коли альфа- і бета-помилки моделювання близькі між собою та загальна точність класифікації набуває максимального значення) є дещо вищою – у середньому на рівні 0,56. Тобто, у разі оптимізації рівня поділу класів ці самі моделі демонстрували б вищі показники ефективності. Однак, у даному дослідженні не було завдання максимального підвищення точності класифікації окремими моделями – необхідно дослідити можливість збільшення ефективності моделювання за рахунок поєднання моделей в ансамблі. І, враховуючи, що зміщення границі поділу між класами має систематичний характер, то висновки, отримані для таких комітетів, будуть релевантними і при їх утворенні з моделей із оптимізованими рівнями поділу класів.

Зауважимо, що база даних, використана у попередніх дослідженнях, і сформований у праці [11] набір показників застосовуватиметься і в цій роботі з метою забезпечення порівнянності результатів моделювання та розвитком авторського підходу до створення ансамблів скорингових моделей оцінювання кредитних ризиків позичальників-фізичних осіб.

V. ЗАСТОСУВАННЯ АЛГОРИТМУ БУСТІНГУ

Як зазначалось вище, однією з умов підвищення ефективності роботи комітету є незалежність похибок його окремих моделей.

Для забезпечення виконання цієї умови розроблено алгоритм бустінгу (підсилення). В його основі лежить ідея послідовного навчання експертів, кожен з яких не зможе повторити помилки попереднього, оскільки буде навчатись на іншому масиві даних. Однаковим для всіх експертів є лише обсяг навчальної вибірки.

Для реалізації бустінгу використовувалась схема алгоритму, запропонована у праці В. Царегородцева [13]. Процедура застосування бустінгу полягає в послідовному навчанні обраної базової моделі з метою підвищення її якості. Базовою називається модель певної конфігурації, що реалізується у трьох варіантах, які відрізняються між собою за рахунок оптимізації на різних навчальних вибірках. Отримані варіанти базової моделі називаються експертами. Процедура формування навчальних вибірок і навчання експертів проходять послідовно.

Позначимо розмірність загального масиву початкових даних M . Для першого експерта з масиву даних M обирається випадковим чином вибірка розмірності N . На отриманих даних проводиться навчання базової моделі, яка і є втіленням першого експерта. За прикладами, які не потрапили до першої навчальної вибірки, здійснюється розрахунок імовірності дефолту першим експертом, за результатами якої формується навчальна вибірка для другого експерта. Ця вибірка також матиме розмірність N і складатиметься наполовину з елементів, які були правильно класифіковані першим експертом, а на другу половину – з елементів, які були класифіковані першим експертом некоректно.

Після навчання другого експерта дані, які не використовувались для формування двох перших навчальних вибірок, використовуються для формування третьої навчальної вибірки. Це здійснюється таким чином. Перші два експерти проводять класифікацію нових прикладів і до третьої навчальної вибірки включають ті N прикладів, за якими два перші експерти дають протилежні результати класифікації. Як можемо бачити з описаної процедури, алгоритм бустінгу, викладений у праці [3], полягає у «послідовній фільтрації прикладів попередніми класифікаторами таким чином, що задача для кожного наступного класифікатора стає складнішою». Після навчання третього експерта комітет можна використовувати для проведення класифікації нових прикладів, узагальнюючи результати всіх експертів простим голосуванням.

З огляду на описану вище процедуру формування навчальних вибірок їх обсяг N не може встановлюватись за класичними підхо-

дами поділу даних для навчання та тестування моделей. Наприклад, неможливо під навчальну вибірку виділити 70 % від загального масиву даних, залишивши 30 % для тестування, оскільки після навчання кожного експерта його навчальна вибірка взагалі вилучається з дослідження. Крім того, слід враховувати необхідність отримати потрібну кількість даних з розбіжностями в класифікації першими двома експертами для формування навчальної вибірки третього експерта. Оскільки всі експерти представлені моделями одного типу, то, як правило, розбіжностей у класифікації на невеликих масивах даних є небагато. Тому обсяг навчальної вибірки для невеликих масивів початкових даних у кожному окремому випадку встановлюється емпірично.

У нашому дослідженні повна база даних складалась із 2075 спостережень. Для вибору максимально можливого обсягу навчальної вибірки було проведено ряд експериментальних розрахунків. Реалізація алгоритму здійснювалась для навчальних вибірок у чотирьох варіантах: 100, 300, 400 і 500 спостережень. В останньому випадку (500 кредитів у навчальній вибірці) було вичерпано всі значення з розбіжностями в класифікації перших двох експертів, тобто збільшити навчальну вибірку понад 500 спостережень було неможливо. Отже, для наявного початкового масиву даних навчальна вибірка може містити не більше 24 %.

Деталізовану схему алгоритму бустінгу для досліджуваних даних проілюстровано на рис. 5.

За базову модель для всіх варіантів розрахунків було обрано багат шарову нейромережу, оскільки в попередніх дослідженнях моделі такої архітектури демонстрували вищі показники ефективності на наявних даних. Питання вибору типу та конфігурації нейромережі для кожного варіанту навчальної вибірки вирішувалось у результаті проведення експериментальних розрахунків. Розглядалось кілька варіантів конфігурацій нейромереж, які демонстрували найвищі показники точності класифікації. Перевага надавалась моделі, яка мала меншу кількість нейронів на прихованому шарі та показувала вищу точність класифікації на тестовому масиві даних. Було виявлено, що при невеликих розмірах навчальних вибірок доцільно використовувати радіально-базисні мережі з меншою кількістю нейронів проміжного шару. В табл. 5 наведено показники ефективності трьох конфігурацій нейромереж, які продемонстрували найвищу точність класифікації для навчальної вибірки з 500 спостережень.

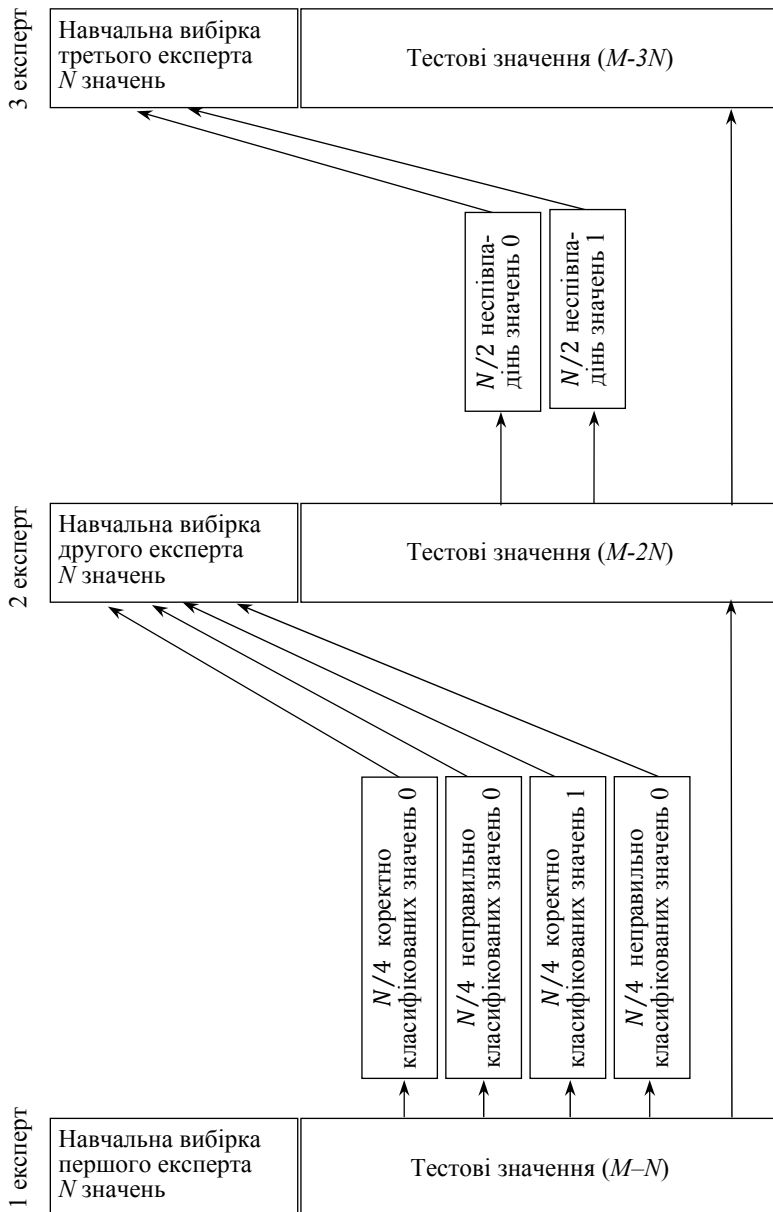


Рис. 5. Схема алгоритму бустінгу для навчальної вибірки з N значень

Таблиця 5

**РОЗРАХУНКИ ЕФЕКТИВНОСТІ НЕЙРОМЕРЕЖ ПРИ ВИБОРІ ПАРАМЕТРІВ
БАЗОВОЇ МОДЕЛІ ДЛЯ НАВЧАЛЬНОЇ ВИБІРКИ 500 СПОСТЕРЕЖЕНЬ**

Архітектура мережі, кількість входів, кількість нейронів проміжного шару	Відсоток правильно класифікованих спос- режень у навчальній вибірці, %	Відсоток правильно класифікованих спос- режень у тестовій ви- бірці, %	Узагальнені дані по всьому масиву (без поділу на навчальну та тестову вибірку)		Правильно класифіко- ваних, %
			Клас	Всього	
Радіально-базисна, 6, 33	75,6	52,8	0	1131	55,1
			1	1044	61,3
Радіально-базисна, 6, 50	77,4	51,4	0	1131	54,5
			1	1044	60,5
Радіально-базисна, 6, 76	79,2	52,0	0	1131	55,3
			1	1044	61,5

На основі даних, наведених у табл. 5, можна зробити висновок, що обраний обсяг навчальної вибірки не надає можливості якісного навчання моделі. Про це свідчить значна розбіжність між показниками точності класифікації для навчальної (понад 75 %) та тестової (не перевищує 53 %) вибірок. Тобто модель, налаштована на невеликій навчальній вибірці, не узагальнює в достатній мірі закономірності поведінки позичальників для ефективного моделювання усього різноманіття варіантів із тестової вибірки. Для підвищення точності класифікації доцільно було б збільшити обсяг навчальної вибірки. Однак, як уже наголошувалось, у цьому дослідженні такої можливості не було через незначний обсяг наявних даних.

На основі даних, викладених у табл. 5, базовою моделлю для навчальної вибірки з 500 спостережень було обрано радіально-базисну нейромережу з 33 нейронами проміжного шару. Така мережа є кращою за показниками ефективності на тестовій вибірці та за узагальненими даними по всьому масиву наявних прикладів.

Аналогічним чином було обрано конфігурацію мереж, які реалізують першого експерта, для трьох інших варіантів навчальних вибірок. Параметри обраних нейромереж і показники їх ефективності наведено в табл. 6.

Таблиця 6

**ПАРАМЕТРИ НЕЙРОМЕРЕЖ ТА ПОКАЗНИКИ ЕФЕКТИВНОСТІ
БАЗОВИХ МОДЕЛЕЙ ПЕРШОГО ЕКСПЕРТА
ДЛЯ ТРЬОХ ВАРІАНТІВ НАВЧАЛЬНОЇ ВИБІРКИ**

Обсяг навчальної вибірки	Архітектура мережі, кількість входів, кількість нейронів проміжного шару	Відсоток правильно класифікованих спостережень у навчальній вибірці, %	Відсоток правильно класифікованих спостережень у тестовій вибірці, %
100	Радіально-базисна, 6, 25	69,0	51,5
300	Радіально-базисна, 6, 35	82,7	52,8
400	Радіально-базисна, 6, 38	76,8	53,9
500	Радіально-базисна, 6, 33	75,6	52,8

Аналіз даних, наведених у табл. 6, вказує на значну розбіжність між показниками правильно класифікованих спостережень на навчальній і тестовій вибірці. Збільшення обсягу навчальної вибірки від 100 до 500 прикладів майже не вплинуло на ефективність класифікації на тестових даних. У даному випадку всі варіанти формування навчальної вибірки не дають змоги здійснити побудову ефективної базової моделі.

Подальша процедура реалізації алгоритму бустінгу використовує обрані базові моделі на нових навчальних вибірках для формування двох наступних експертів ансамблю. Зміни ефективності нових експертів для кожного варіанту базової моделі, що відповідають різним обсягам навчальної вибірки, наведено в табл. 7.

Завдання класифікації для наступних двох експертів є складнішим, ніж для першого. Тому всі базові моделі демонструють зниження точності класифікації на навчальній вибірці, як можемо бачити в табл. 7. Проте на тестових даних спостерігається зростання ефективності моделювання. Отже, в процесі реалізації алгоритму бустінгу навчальна вибірка стає різноманітнішою та набуває прикладів з відмінними між собою характеристиками

(у тому числі такими, що погано ідентифікувались попередніми експертами), що приводить до ефективнішого узагальнення властивостей і покращення здатності експертів розрізняти класи на нових даних.

Таблиця 7

**ЗМІНИ ЕФЕКТИВНОСТІ БАЗОВИХ МОДЕЛЕЙ ДЛЯ ДРУГОГО
ТА ТРЕТЬОГО ЕКСПЕРТІВ ПРИ ТРЬОХ ВАРІАНТАХ НАВЧАЛЬНОЇ ВИБІРКИ**

Архітектура мережі, кількість входів, кількість нейронів проміжного шару	Обсяг навчальної вибірки	Другий експерт			Третій експерт		
		Відсоток правильно класифікованих спостережень у навчальній вибірці, %	Відсоток правильно класифікованих спостережень у тестовій вибірці, %	Обсяг тестових вибірок	Відсоток правильно класифікованих спостережень у навчальній вибірці, %	Відсоток правильно класифікованих спостережень у тестовій вибірці, %	Обсяг тестових вибірок
Радіально-базисна, 6, 25	100	62,0	55,5	1975	62,0	56,6	1875
Радіально-базисна, 6, 35	300	61,7	54,4	1575	58,3	57,2	1275
Радіально-базисна, 6, 38	400	57,3	55,5	1375	56,5	58,9	975
Радіально-базисна, 6, 33	500	62,0	55,1	1175	55,8	58,8	675

Алгоритм бустінгу було винайдено для можливостей підсилення слабких моделей. Отже, показником ефективності застосування даного алгоритму є порівняння точності класифікації трьох експертів (базової моделі, побудованої на трьох різних вибірках) та комітету, сформованого на їх основі. Оскільки всі приклади, що входили до навчальних вибірок, вилучались із подальших досліджень, то порівняння ефективності усіх експертів та комітету здійснювалось лише на даних останньої тестової вибірки для своєї базової моделі. Відповідні розрахунки показників ефективності у розрізі класів надійних і ненадійних позичальників наведено в табл. 8.

Таблиця 8
ПОРІВНЯННЯ ЕФЕКТИВНОСТІ ПЕРШОГО, ДРУГОГО, ТРЕТЬОГО ЕКСПЕРТІВ І КОМІТЕТУ МОДЕЛЕЙ
ПРИ ТРЬОХ ВАРІАНТАХ НАВЧАЛЬНОЇ ТА ТЕСТОВОЇ ВИБІРОК

	Відсоток правильно класифікованих спостережень першим експертом, %		Відсоток правильно класифікованих спостережень другим експертом, %		Відсоток правильно класифікованих спостережень третім експертом, %		Відсоток правильно класифікованих спостережень комітетом експертів, %	
	надійні позичальники	ненадійні позичальники	надійні позичальники	ненадійні позичальники	надійні позичальники	ненадійні позичальники	надійні позичальники	ненадійні позичальники
Базова модель								
Радіально-базисна, 6, 25 (навчальна вибірка – 100, тестова вибірка – 1875 спостережень)	49,3	54,2	53,3	58,4	53,6	60	53,6	60,8
Загальний відсоток правильно класифікованих спостережень %	51,7		55,8		56,6		57,1	
Радіально-базисна, 6, 35 (навчальна вибірка – 300, тестова вибірка – 1275 спостережень)	50,5	58,2	53,4	57,6	54,2	60,1	52,7	62,1
Загальний відсоток правильно класифікованих спостережень %	54,1		55,4		57,2		57,1	

Закінчення табл. 8

Базова модель	Відсоток правильно класифікованих спостережень першим експертом, %		Відсоток правильно класифікованих спостережень другим експертом, %		Відсоток правильно класифікованих спостережень третім експертом, %		Відсоток правильно класифікованих спостережень комітетом експертів, %	
	надійні позичальники	ненадійні позичальники	надійні позичальники	ненадійні позичальники	надійні позичальники	ненадійні позичальники	надійні позичальники	ненадійні позичальники
Радіально-базисна, 6, 38 (навчальна вибірка – 400, тестова вибірка – 975 спостережень)	49,5	66,2	55,2	60,8	56,1	62,2	53,3	63,1
Загальний відсоток правильно класифікованих спостережень %	57,1		57,7		58,8		57,7	
Радіально-базисна, 6, 33 (навчальна вибірка – 500, тестова вибірка – 675 спостережень)	45,4	71,2	53,4	70,0	45,9	77,7	51,9	70,4
Загальний відсоток правильно класифікованих спостережень %	55,9		60,1		58,8		59,4	

Як можна бачити з табл. 8, майже у всіх випадках експерти після перенавчання на нових вибірках підвищують точність класифікації. Виключення складає лише третій експерт за базовою моделлю із навчальною вибіркою з 500 спостережень. Таку невідповідність можна пояснити ще суттєвішим зміщенням границі раціонального поділу між класами від встановленого у пакеті STATISTICA рівня 0,5, що можна бачити за значно збільшеної різниці між альфа- та бета-помилками моделювання, яка, в свою чергу, впливає на зниження загальної точності класифікації. Крім того, на відміну від попередніх випадків, при формуванні навчальної вибірки для цього експерта було включено всі приклади із розбіжностями в моделюванні ненадійних позичальників першими двома експертами. Відповідно, у тестовій вибірці для перевірки адекватності четвертого комітету моделей не залишається прикладів ненадійних позичальників для здійснення розрахунків за третім експертом. Через це висока ефективність ідентифікації цією моделлю ненадійних клієнтів не враховується при їх розпізнаванні комітетом (саме тому в табл. 8 точність класифікації ненадійних позичальників комітетом за четвертою базовою моделлю знаходиться посередині між першими двома експертами), що дещо знижує загальний показник ефективності цього комітету. Однак, навіть незважаючи на це, всі чотири варіанти утворених комітетів демонструють підсилення ефективності базових моделей – точність класифікації комітетом або перевищує кожного з експертів за відповідною базовою моделлю, або є близькою до найкращої з них.

Різновиди алгоритмів бустінгу (наприклад, адаптивного бустінгу AdaBoost) надають можливість створювати ефективніші комітети моделей. В AdaBoost використовується спеціальна процедура збільшення в навчальних вибірках кількості прикладів, на яких були допущені помилки попередніми експертами. Також при узагальненні результатів за цим алгоритмі враховуються відносні ваги експертів комітету, які визначаються залежно від їх помилок (вага зменшується при збільшенні похибки класифікації). Для такого варіанту алгоритму Р. Шапайра [14] доведено, що похибка комітету буде зменшуватись експоненційно, якщо похибки окремих експертів менші 0,5.

VI. АНСАМБЛЬ НА ОСНОВІ СПЕЦІАЛІЗАЦІЇ ЕКСПЕРТІВ

При моделюванні кредитних ризиків важливо враховувати суттєву неоднорідність досліджуваних даних. Так, клієнти креди-

тних установ мають значні відмінності за показниками, які використовуються для оцінювання кредитоспроможності. Наприклад, вікова категорія позичальників може змінюватись в межах від 20 до 60 і більше років, рівень освіти – від середньої до наявності двох чи більше вищих освіт тощо. І різні категорії позичальників характеризуються різним рівнем кредитного ризику. Тому виділення з усього наявного масиву даних більш однорідної вибірки, наприклад із позичальників лише з вищою освітою, дасть змогу досліджувати таку їх групу, яка характеризується більш-менш подібними соціально-економічними умовами існування, що дозволить ефективніше виявляти закономірності поведінки позичальників.

Для обробки таких даних доцільно застосовувати підходи до утворення ансамблів, що враховують спеціалізацію експертів. Такі ансамблі відносяться до категорії динамічного об'єднання моделей. У С. Терехова [3] було описано одну з ідей створення ансамблю даної категорії, яка базується на застосуванні карт Кохонена. Такий підхід вимагає додаткових досліджень щодо вибору параметрів карти, які впливатимуть на результат кластеризації. Універсальних рекомендацій з даного питання немає, тож розмір карт самоорганізації визначається індивідуально для кожної задачі, що обґрунтовано у праці А. Матвійчука [15]. Крім того, поділ масиву даних на кластери призведе до ще більшого зменшення обсягів навчальних вибірок, що дасть негативний вплив на результати класифікації та може зменшити ефективність моделювання (у випадку невеликої початкової вибірки даних).

Однак реалізація ідеї спеціалізації експертів комітету є дуже привабливою для розв'язання задачі моделювання кредитних ризиків, тому для її втілення було вирішено випробувати підхід, за яким поділ наявних даних на окремі групи не буде базуватись на процедурі кластеризації.

Пропонується такий алгоритм формування ансамблю моделей. З масиву початкових даних виділяється K якісних показників, які будуть використані при оцінюванні кредитоспроможності позичальників. Для кожного якісного показника визначається кількість його категорій N_j , $j = \overline{1, K}$, за якими із загального масиву спостережень виділяються підгрупи, які складатимуть навчальні вибірки для відповідних моделей-експертів

$L_{ij}, i = \overline{1, N_j}, j = \overline{1, K}$. Подібне формування вибірок для оптимізації окремих моделей робить можливою реалізацію ідеї спеціалізації експертів. За потреби, кожен експерт може бути представлений моделями різних типів – логістичні регресії, нейромережі, дерева рішень тощо.

Загальний результат роботи комітету можна визначати трьома способами: простим чи зваженим голосуванням або за допомогою метамоделі. Для випадку простого голосування використовуються рейтинги моделей-експертів $R_{ij}, i = \overline{1, N_j}, j = \overline{1, K}$. Ці рейтинги встановлюються на основі рівня значущості (зокрема, для логістичних регресій) або за рівнем точності класифікації, що демонструє модель на тестових даних (для випадку застосування нейромереж).

Для аналізу кредитоспроможності нового потенційного клієнта обирається група експертів M , які можуть бути задіяні для його класифікації (для застосування яких за цим позичальником є всі необхідні дані та навчальні вибірки яких формувались на значеннях якісних показників, відповідних даному клієнту). Отже, оцінювання кредитного ризику цього позичальника здійснюється комітетом, утвореним з моделей $L_{ij} \in M$. Для проведення розрахунків моделі впорядковуються за їх рейтингами R_{ij} . Аналіз позичальника починається з моделей, які мають найвищий рейтинг. Коли більше половини задіяних експертів віднесли аналізованого позичальника до одного класу, процес оцінювання кредитоспроможності припиняється. Якщо виникає ситуація, коли кількість експертів є парною і голоси розділяються порівну, то процедуру узагальнення результатів слід замінити на зважене голосування.

У випадку зваженого голосування для кожного експерта фіксується відсоток правильно класифікованих на тестовій вибірці прикладів $t_{ij}, i = \overline{1, N_j}, j = \overline{1, K}$. При обробці даних потенційного клієнта визначається група відповідних йому експертів M . На основі значень t_{ij} обчислюються коефіцієнти компетентності для обраних експертів $L_{ij} \in M$:

$$k_{ij} = \frac{t_{ij}}{\sum_{L_{ij} \in M} t_{ij}}.$$

Загальний результат розрахунків ансамблю є сумою добутоків виходів окремих моделей на їх коефіцієнти компетентності. Основою для визначення коефіцієнтів компетентності експерта може також слугувати будь-який інший показник точності моделі, наприклад коефіцієнт Джині.

Третім способом узагальнення результатів роботи комітету є використання мета моделі – моделі верхнього рівня, входами якої будуть розрахунки окремих експертів.

Схематично зображення описаного вище алгоритму побудови ансамблю моделей наведено на рис. 6.

Так, в авторському дослідженні [11] для оцінювання рівня кредитоспроможності фізичної особи запропоновану методику побудови ансамблю реалізовано для випадку використання *logit*-моделей. Враховано вплив таких якісних показників, як рівень освіти та статус працюючого. За кожним з цих показників із початкового масиву даних утворюється по кілька підгруп. Наприклад, показник «рівень освіти» поділяється на два підрівні: «наявність однієї чи більше вищих освіт» і «наявність середньої та середньої спеціальної освіти». З усього масиву даних обираються лише ті записи, які відповідають позичальникам із певним підрівнем даного якісного показника, утворюючи таким чином два окремі масиви більш однорідних даних. Аналогічна процедура здійснюється за іншим показником (статусом працюючого). Додатково було здійснено поділ початкової вибірки на три підгрупи за такою важливою з точки зору оцінки кредитоспроможності фізичних осіб характеристикою, що визначає наявність утриманців. Сформовані таким чином масиви даних реалізують ідею спеціалізації експертів.

За запропонованим підходом початковий масив розбивається на підгрупи, які перетинаються між собою. У такому вигляді формування масивів даних для навчання окремих моделей-експертів схоже на відповідну процедуру за технологією беггінгу. Однак за алгоритмом беггінгу вибірки утворюються випадковим чином, а в запропонованому варіанті формування вибірок є логічно обумовленим і приводить до отримання більш однорідних масивів даних.

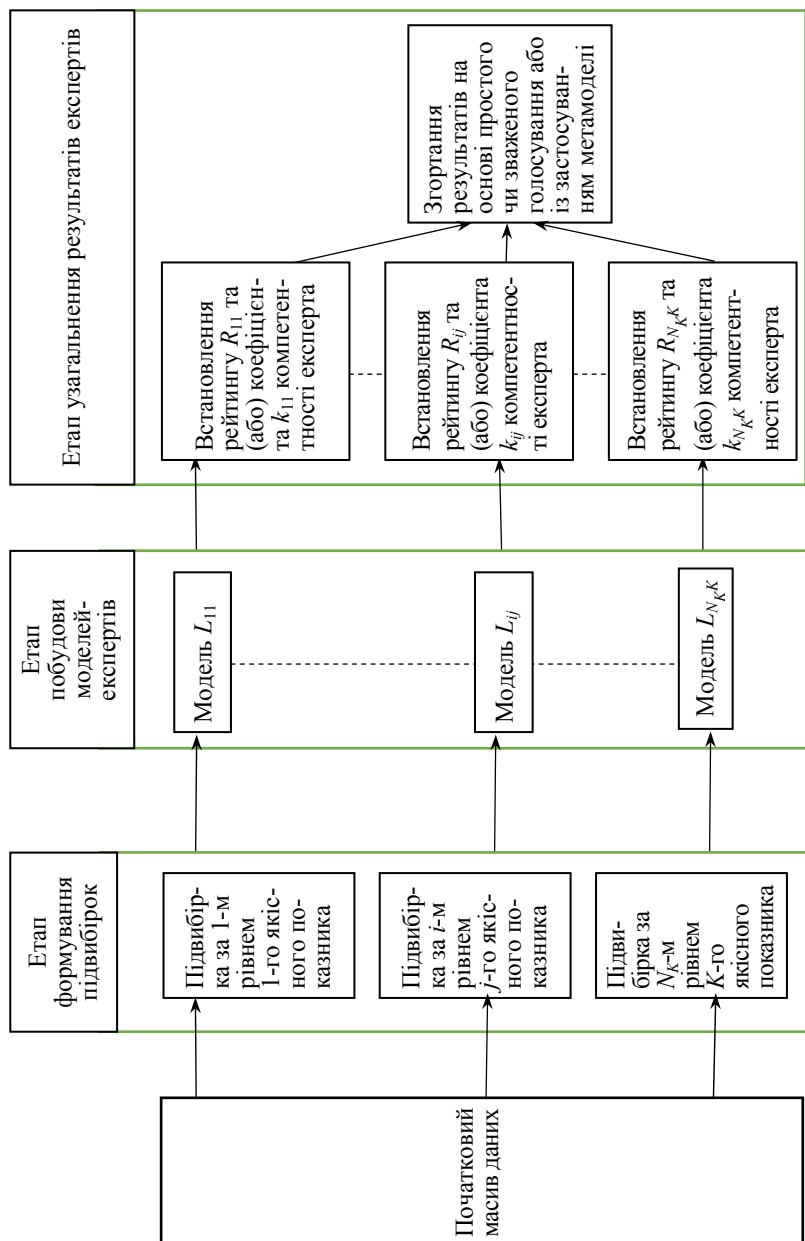


Рис. 6. Схематичне зображення процесу побудови ансамблю моделей на основі спеціалізації експертів

Очевидно, що можна будувати моделі, використовуючи будь-яку потрібну кількість якісних показників. Узагальнення результатів розрахунків усіх моделей доцільно здійснювати шляхом простого або зваженого голосування. Ваги для моделей відповідатимуть рівню точності класифікації експерта (чим він є вищим, тим більшою є його вага).

Сформована таким чином технологія побудови ансамблю найкращим чином пристосована саме для специфіки задачі оцінювання кредитоспроможності позичальників – фізичних осіб. Вона поєднує в собі й ідею спеціалізації експертів, і логічний спосіб формування масивів даних для навчальних вибірок, і класичний спосіб поєднання результатів розрахунків окремих моделей ансамблю.

Варто зазначити, що за реалізації описаного підходу також не вдається уникнути складностей, пов'язаних із зазначеною вище проблемою недостатньо великого обсягу наявних даних. Крім того, утворені підгрупи мають значні розбіжності в обсягах: найменша містить 145 спостережень, найбільша – 1547. До того ж, за кожною з них потрібно утворити навчальну та тестову вибірки. Зауважимо, що при визначенні розміру навчальної вибірки недоцільно орієнтуватись на найменший масив даних, оскільки тоді втрачається можливість кращого налаштування інших моделей на більших вибірках. Таким чином виникає проблема неможливості формування навчальних вибірок однакового розміру. А отже, не вдається повністю коректно застосувати запропоновану схему алгоритму побудови ансамблю моделей.

Враховуючи особливості розподілу позичальників між класами та підгрупами в досліджуваній базі даних, у кожній підгрупі виділялось близько 80 % записів для навчальної вибірки та 20 % – для тестової. Для вибору експертів за кожною підгрупою розглядалися три кращі мережі за показниками точності класифікації на навчальній і тестовій вибірках, а також по узагальненим даним. Результати проведених обчислень наведено в табл. 9. Найефективніші моделі за кожною підгрупою тут виділені напівжирним шрифтом.

Аналіз даних, наведених у табл. 9, підтверджує отримані вище висновки, що ефективнішу класифікацію за представленими даними здійснюють мережі з радіально-базисною архітектурою, і лише в окремих випадках кращий результат демонструє тришаровий персептрон (він виявився ефективнішим для масивів меншої розмірності).

Таблиця 9

**ПОКАЗНИКИ ЕФЕКТИВНОСТІ НЕЙРОМЕРЕЖ РІЗНОЇ АРХІТЕКТУРИ
ДЛЯ ВИБОРУ ОКРЕМИХ МОДЕЛЕЙ-ЕКСПЕРТІВ
ПРИ ФОРМУВАННІ АНСАМБЛЮ**

Архітектура мережі, кількість входів, кількість нейронів проміжного шару	Відсоток правильно класифікованих спостережень у навчальній вибірці, %	Відсоток правильно класифікованих спостережень у тестовій вибірці, %	Узагальнені дані по всьому масиву (без поділу на навчальну та тестову вибірки)		Правильно класифікованих, %
			Клас	Всього	
Підгрупа позичальників із «наявністю однієї чи більше вищих освіт», відібраних за показником «рівень освіти»					
Тришаровий персептрон, 6, 6	64	58	0	439	61,73
			1	297	60,27
Радіально-базисна, 6, 9	59	53	0	439	56,04
			1	297	55,89
Радіально-базисна, 6, 4	57	53	0	439	55,35
			1	297	55,22
Підгрупа позичальників із «наявністю середньої та спеціальної середньої освіти», відібраних за показником «рівень освіти»					
Тришаровий персептрон, 6, 7	62	61	0	692	63,15
			1	747	60,1
Радіально-базисна, 6, 7	62	61	0	692	63,29
			1	747	60,37
Радіально-базисна, 6, 30	65	58	0	692	63
			1	747	59,97
Підгрупа позичальників із «власною справою», відібраних за показником «статус працівника»					
Радіально-базисна, 6, 3	60	56	0	76	60,53
			1	69	55,07
Радіально-базисна, 6, 6	63	58	0	76	64,47
			1	69	56,52

Продовження табл. 9

Архітектура мережі, кількість входів, кількість нейронів проміжного шару	Відсоток правильно класифікованих спостережень у навчальній вибірці, %	Відсоток правильно класифікованих спостережень у тестовій вибірці, %	Узагальнені дані по всьому масиву (без поділу на навчальну та тестову вибірки)		Правильно класифікованих, %
			Клас	Всього	
Радіально-базисна, 6, 4	58	58	0	76	60,53
			1	69	55,07
Підгрупа «найманих працівників», відібраних за показником «статус працівника»					
Радіально-базисна, 6, 4	60	58	0	809	58,47
			1	738	60,43
Радіально-базисна, 6, 16	61	59	0	809	58,96
			1	738	60,98
Радіально-базисна, 6, 8	60	59	0	809	58,96
			1	738	60,98
Підгрупа позичальників із «іншим статусом», відібраних за показником «статус працівника»					
Тришаровий персептрон, 6, 9	64	50	0	246	56,5
			1	237	57,8
Радіально-базисна, 6, 7	61	53	0	246	57,8
			1	237	57,32
Радіально-базисна, 6, 3	60	56	0	246	58,65
			1	237	54,88
Підгрупа даних «відсутні», відібраних за показником «наявність утриманців»					
Радіально-базисна 5, 11	60	59	0	491	59,06
			1	489	60,33

Закінчення табл. 9

Архітектура мережі, кількість входів, кількість нейронів проміжного шару	Відсоток правильно класифікованих спостережень у навчальній вибірці, %	Відсоток правильно класифікованих спостережень у тестовій вибірці, %	Узагальнені дані по всьому масиву (без поділу на навчальну та тестову вибірки)		Правильно класифікованих, %
			Клас	Всього	
Радіально-базисна, 5, 15	63	59	0	491	60,29
			1	489	61,55
Радіально-базисна, 6, 23	63	58	0	491	60,29
			1	489	61,55
Підгрупа даних «одна особа», відібраних за показником «наявність утриманців»					
Тришаровий персептрон, 3, 8	56	58	0	435	57,01
			1	355	57,46
Радіально-базисна, 5, 2	57	55	0	435	56,32
			1	355	56,06
Радіально-базисна, 5, 5	59	58	0	435	57,93
			1	355	58,59
Підгрупа даних «дві та більше особи», відібраних за показником «наявність утриманців»					
Радіально-базисна, 5, 3	49	54	0	205	54,15
			1	200	49
Радіально-базисна, 5, 6	68	61	0	205	68,29
			1	200	60,5
Радіально-базисна, 5, 8	67	60	0	205	66,83
			1	200	59,5

Результати розрахунків усіх моделей ансамблю узагальнювались простим голосуванням. Висновок про точність класифікації,

що проведена ансамблем моделей за запропонованою технологією, можна зробити на основі перевірки за усім початковим масивом даних, проводячи порівняння отриманого результату з результатами розрахунків інших комітетів або окремих нейромереж.

За всім масивом даних комітетом було правильно класифіковано 63,5 % надійних позичальників та 60 % дефолтних. Відповідні показники для ансамблю моделей, сформованих за алгоритмом бустінгу, складали 58,7 % правильно класифікованих надійних позичальників і 56 % – дефолтних. Як можна бачити, для досліджуваних даних кращі результати демонструє ансамбль, який враховує спеціалізацію експертів.

Метою створення алгоритму бустінгу було підвищення точності класифікації у великих масивах інформації, яких при проведенні дослідження у нас в наявності не було. Також цей алгоритм не враховує особливостей даних, на яких проводяться розрахунки. Оскільки задача класифікації позичальників банку значною мірою пов'язана зі специфікою початкових даних, наприклад, наявністю значної кількості якісних показників, то від коректного врахування особливостей даних залежить і точність класифікації. Отже, при розв'язанні поставленої задачі доцільно застосовувати такий підхід для формування ансамблю, в якому враховується спеціалізація окремих моделей-експертів. Крім того, запропонований підхід може бути застосований для випадків, коли початковий масив даних має невелику розмірність і більшість алгоритмів формування ансамблів не дають задовільних результатів. Такі задачі виникатимуть, зокрема, при дослідженнях нових банківських продуктів, які ще не мають широкого розповсюдження, а, отже, обсяги статистичної інформації за ними недостатні для застосування класичних технологій формування ансамблів.

Оскільки реалізований алгоритм побудови ансамблю було також застосовано для випадку використання в якості експертів *logit*-моделей, то проведемо порівняння ефективності обох цих комітетів. Однак безпосередньо використовувати для порівняння розрахунки, описані в праці [11], неможливо, оскільки побудова *logit*-регресій здійснювалась на всіх даних підгрупи і не містила поділу на навчальну та тестову вибірки. Для коректнішого порівняння ефективності комітетів, утворених з *logit*-моделей і нейромереж, по кожній підгрупі проведено розрахунки параметрів

logit-регресій на основі лише тієї частини даних, яка була використана в якості навчальних вибірок для нейронних мереж. У результаті отримано вісім логістичних регресій, показники точності яких наведено в табл. 10.

Таблиця 10

ПОКАЗНИКИ ТОЧНОСТІ *LOGIT*-РЕГРЕСІЙ

Показник, що використовувався для формування підгрупи даних	Навчальна вибірка		Тестова вибірка		Значення χ^2
	Відсоток правильно класифіко- ваних зна- чень «0»	Відсоток правильно класифіко- ваних зна- чень «1»	Відсоток правильно класифіко- ваних зна- чень «0»	Відсоток правильно класифіко- ваних зна- чень «1»	
Рівень освіти – «наявність одної чи більше вищих освіт»	60,8	66,3	31	67	47,561
Рівень освіти – «наявність середньої та спеціальної се- редньої освіти»	59,6	66,8	62	48	98,497
Статус працівника – «власна справа»	60,0	58,2	45	48	16,164
Статус працівника – «найманий праців- ник»	62,5	67,0	45	56	124,110
Статус працівника – «інший статус»	56,3	62,6	56	65	29,389
Наявність утриман- ців – «відсутні»	58,8	66,5	52	30	67,646
Наявність утриман- ців – «одна особа»	58,5	62,0	40	30	37,573
Наявність утриман- ців – «дві та більше особи»	59,4	69,4	11	78	36,753

Аналіз даних, наведених у табл. 10, свідчить, що всі моделі, крім побудованої на даних по статусу працівника «власна справа», є статистично значущими на навчальних вибірках (масив навчальних даних для вказаної моделі є найменшим у дослідженні та складає лише 110 прикладів). Точність класифікації на навчальних вибірках для всіх моделей майже однакова із дещо вищою ефективністю визначення ненадійних позичальників. На тестових вибірках у більшості моделей спостерігається значне зниження точності класифікації. Отримані результати свідчать про недоцільність застосування *logit*-регресій для розв'язання задачі класифікації за відсутності достатньої кількості навчальних даних. Узагальнений результат розрахунків цих моделей за описаним вище алгоритмом побудови ансамблю продемонстрував точність передбачення надійних позичальників на рівні 56,6 %, а дефолтних – 61,8 %.

Підсумовуючи результати експериментальних розрахунків можна зробити висновок, що комітети моделей, які були утворені на основі всього масиву даних (без поділу на підгрупи) або ж на основі використання *logit*-регресій, недоцільно використовувати для розв'язання поставленої задачі оцінювання кредитних ризиків позичальників-фізичних осіб. Комітет моделей, сформований на повному масиві даних, в окремих випадках демонструє точність оцінювання менше 50 %. Так само неприйнятні показники точності класифікації мають окремі моделі *logit*-регресій, які входять у комітет.

Таким чином, за умов неоднорідності в масиві початкових даних доречно використовувати лише два комітети: отриманий за алгоритмом бустінгу та комітет нейромереж з урахуванням спеціалізації моделей-експертів.

Як наголошувалось вище, ефективнішими є комітети, при формуванні яких враховано дві умови – в їх основу покладено моделі, які мають вищі показники ефективності (у даній роботі продемонстровано відповідність зазначеній умові процедури пошуку та відбору кращих моделей при формуванні всіх описаних видів комітетів), а також забезпечується незалежність похибок окремих моделей ансамблю. У праці [3] зазначається, що для перевірки статистичної незалежності розподілу похибок окремих моделей комітету можна скористатись спрощеним підходом аналізу попарних кореляцій похибок усіх пар моделей. З цією метою розглянемо кореляційні матриці похибок моделю-

вання, які утворені двома зазначеними видами комітетів моделей. Для аналізу кореляційної матриці при застосуванні алгоритму бустінгу обрано найкращий варіант ансамблю (навчальна вибірка 500 значень). У цьому випадку кореляційна матриця має вигляд:

$$R_{Boost} = \begin{pmatrix} 1 & 0,88 & 0,64 \\ 0,88 & 1 & 0,65 \\ 0,64 & 0,65 & 1 \end{pmatrix}.$$

Досить тісний зв'язок спостерігається між похибками першого та другого експертів, менш тісний – для інших пар.

Кореляційна матриця похибок для ансамблю на основі спеціалізації експертів відображає всі пари зв'язків восьми експертів і має вигляд:

$$R_{Expert} = \begin{pmatrix} 1 & 0,002 & 0,28 & 0,24 & 0,07 & 0,16 & 0,36 & 0,09 \\ 0,002 & 1 & 0,44 & 0,36 & 0,28 & 0,11 & 0,54 & 0,34 \\ 0,28 & 0,44 & 1 & -0,0003 & 0,0002 & 0,08 & 0,39 & 0,25 \\ 0,24 & 0,36 & -0,0003 & 1 & 0,001 & 0,14 & 0,42 & 0,22 \\ 0,07 & 0,28 & 0,0002 & 0,001 & 1 & 0,08 & 0,21 & 0,15 \\ 0,16 & 0,11 & 0,08 & 0,14 & 0,08 & 1 & 0,0003 & 0,00009 \\ 0,36 & 0,54 & 0,39 & 0,42 & 0,21 & 0,0003 & 1 & -0,0004 \\ 0,09 & 0,34 & 0,25 & 0,22 & 0,15 & 0,00009 & -0,0004 & 1 \end{pmatrix}.$$

Найвище значення коефіцієнту кореляції у матриці для восьми експертів сягає 0,54, що є нижчим, ніж найменше значення кореляції похибок окремих моделей у комітеті, сформованим за алгоритмом бустінгу. Причому, більшість показників кореляції похибок моделей в ансамблі на основі спеціалізації експертів знаходяться у межах від 0 до 0,2. Таким чином, на основі проведеного порівняльного аналізу можна зробити висновок щодо незалежності похибок окремих моделей ансамблю, сформованого за алгоритмом врахування спеціалізації експертів, що свідчить про його вищу ефективність за алгоритм бустінгу [3].

Розглянемо результати класифікації, які отримані двома зазначеними ансамблями моделей. Показники їх ефективності наведено в табл. 11.

Таблиця 11

**ПОКАЗНИКИ ТОЧНОСТІ КЛАСИФІКАЦІЇ АНСАМБЛЯМИ МОДЕЛЕЙ,
СФОРМОВАНИМИ ЗА РІЗНИМИ АЛГОРИТМАМИ**

Види ансамблів	Відсоток правильно класифікованих значень «0»	Відсоток правильно класифікованих значень «1»
Ансамбль, сформований із урахуванням спеціалізації моделей-експертів	63,5	60,0
Ансамбль, сформований з нейромереж за алгоритмом бустінгу (кращий результат з чотирьох ансамблів)	51,9	70,4

Як бачимо за даними табл. 11, запропонована процедура формування ансамблю моделей на основі спеціалізації експертів демонструє рівномірнішу ефективність моделювання в розрізі класів і вищу загальну точність класифікації – 61,8 % (тоді як кращий результат класифікації за обома класами на основі алгоритму бустінгу становив 59,4 %).

VII. ВИСНОВКИ

Узагальнюючи отримані результати можна зробити такі висновки. Найпростішим для реалізації при розробці системи скорингової оцінки позичальників банку є тип ансамблю, робота якого полягає в усередненні результатів розрахунків моделей-експертів. Однак проведене експериментальне дослідження показало, що застосування такого типу ансамблю не дало змоги підвищити ефективність класифікації позичальників – точність ансамблю майже співпадає з точністю кращої з його моделей. Перевагою такого виду ансамблю є підвищення стійкості результату його розрахунків.

Зростання точності класифікації понад точність окремої моделі забезпечує реалізація алгоритму бустінгу. Проте підвищення ефективності моделі є недостатньо високим. Крім того, цей алгоритм рекомендується застосовувати для випадку однорідних початкових даних.

Найбільш адаптованим для проведення скорингової оцінки позичальників-фізичних осіб можна вважати запропонований ав-

торами алгоритм побудови ансамблю на основі спеціалізації експертів. Процедура формування навчальних вибірок для експертів ансамблю забезпечує побудову моделей на однорідних масивах даних, згрупованих за значеннями якісних характеристик позичальників. При цьому окремі експерти комітету можуть бути реалізовані на основі різних типів моделей (логістичні регресії, нейромережі, дерева рішень тощо). В алгоритмі враховується компетентність експертів за рахунок введення вагових коефіцієнтів при узагальненні результатів.

Серед переваг авторського підходу є і те, що більшість відомих ансамблевих технологій вимагають значних обсягів та однорідності початкових даних. Запропонований підхід до утворення ансамблю моделей на основі спеціалізації експертів враховує неоднорідність початкових даних і може бути використаний і для невеликої початкової сукупності спостережень. Такі задачі виникатимуть, наприклад, при дослідженнях нових банківських продуктів, які ще не мають широкого поширення, а, отже, обсяги статистичної інформації за ними недостатні для застосування класичних технологій формування ансамблів.

Проведені в дослідженні експериментальні розрахунки підтвердили переваги запропонованого авторами алгоритму побудови ансамблю на основі спеціалізації експертів при розв'язанні задач оцінювання кредитного ризику.

Список літератури

1. *Miller G. A.* The Magic Number Seven Plus or Minus Two: Some Limits on Our Capacity for Processing Information // *Psychological Review*. – 1956. – № 63. – P. 81–97.
2. *Flach P. A.* Machine Learning: The Art and Science of Algorithms that Make Sense of Data. – Cambridge: Cambridge University Press, 2012. – 409 p.
3. *Терехов С. А.* Гениальные комитеты умных машин // IX Всероссийская научно-техническая конференция «Нейроинформатика-2007»: Лекции по нейроинформатике. — М.: МИФИ, 2007. — Ч. 2. — С. 11–42.
4. *Schapire R. E.* The Strength of Weak Learnability // *Machine Learning*. – 1990. – Vol. 5. – Is. 2. – P. 197–227. – Режим доступу: <http://rob.schapire.net/papers/strengthofweak.pdf>.
5. *Freund Y., Schapire R. E.* A decision-theoretic generalization of on-line learning and an application to boosting // *Journal of Computer and System Sciences*. – 1997. – No. 55. – Vol. 1. – P. 119–139. – Режим доступу: http://www.face-rec.org/algorithms/Boosting-Ensemble/decision-theoretic_generalization.pdf.

6. *Breiman L.* Random forest [Електронний ресурс]. – Berkeley, CA: Statistics Department of University of California, 2001. – 32 p. – Режим доступу: <https://www.stat.berkeley.edu/~breiman/randomforest2001.pdf>.

7. *Кузнецов И. А., Киреев В. С.* Разработка ансамбля алгоритмов классификации с использованием энтропийного показателя качества для решения задачи поведенческого скоринга [Електронний ресурс] // Труды XVIII Международной конференции «Аналитика и управление данными в областях с интенсивным использованием данных» (DAMDID/RCDL'2016). — Ершово, 2016. — С. 11—42. — Режим доступу: <http://ceur-ws.org/Vol-1752/paper07.pdf>.

8. *Паклин Н. Б.* Бизнес-аналитика: от данных к знаниям / Н. Б. Паклин, В. И. Орешков. — СПб.: Питер, 2013. — 704 с.

9. *Хайкин С.* Нейронные сети: полный курс, 2-е издание: Пер. с англ. / С. Хайкин. — М.: Издательский дом «Вильямс», 2006. — 1104 с.

10. *Лавренков Ю. Н.* Исследование и разработка комбинированных нейросетевых технологий для повышения эффективности безопасной маршрутизации информации в сетях связи: дисс... канд. тех. наук: 05.13.17 – «Теоретические основы информатики». — Калуга, 2014. — 208 с. — Режим доступу: https://mpei.ru/Science/Dissertations/dissertations/Dissertations/LavrenkovYN_diss.pdf.

11. *Савіна С. С.* Об'єднання моделей logit-регресій як комітету експертів для оцінки кредитоспроможності позичальника / С. С. Савіна, В. П. Бень // Нейро-нечіткі технології моделювання в економіці. — 2015. — № 4. — С. 154—188.

12. *Савіна С. С.* Вибір архітектури нейромережі для розв'язання задачі класифікації надійності позичальників-фізичних осіб / С. С. Савіна, В. П. Бень // Нейро-нечіткі технології моделювання в економіці. — 2016. — № 5. — С. 123—151.

13. *Царегородцев В. Г.* Оптимизация экспертов boosting-коллектива по их кривым обучения / В. Г. Царегородцев // Материалы XIII Всероссийского семинара «Нейроинформатика и ее приложения». — Красноярск, 2004. — С. 152—157.

14. *Schapire R. E.* Theoretical views of boosting and applications [Електронний ресурс] // Proceedings of 10th International Conference «Algorithmic Learning Theory» (ALT '99). — Tokyo, 1999. — Режим доступу: http://www-ai.cs.uni-dortmund.de/LEHRE/PG/PG445/literatur/schapire_99a.pdf.

15. *Мамсійчук А. В.* Штучний інтелект в економіці: нейронні мережі, нечітка логіка: Монографія. — К.: КНЕУ, 2011. — 439 с.

References

1. Miller, G. A. (1956). The Magic Number Seven Plus or Minus Two: Some Limits on Our Capacity for Processing Information. *Psychological Review*, 63, 81–97.

2. Flach, P. A. (2012). *Machine Learning: The Art and Science of Algorithms that Make Sense of Data*. Cambridge York, UK: Cambridge University Press.
3. Terekhov, S. A. (2007). Genialniye komitety umnykh mashyn. *IX Vserossiyskaya nauchno-tehnicheskaya konferenciya «Nejroinformatika-2007»: Lekcii po nejroinformatike (Moscow, Russia: MIFI), (IX All-Russian scientific and technical conference «Neuroinformatics-2007»: Lectures on neuroinformatics)*, 2, 11—42 [in Russian].
4. Schapire, R. E. (1990). The Strength of Weak Learnability. *Machine Learning*, 5(2), 197—227. Retrieved from <http://rob.schapire.net/papers/strengthofweak.pdf>.
5. Freund, Y., & Schapire, R. E. (1997). A decision-theoretic generalization of on-line learning and an application to boosting. *Journal of Computer and System Sciences*, 55(1), 119—139. Retrieved from http://www.face-rec.org/algorithms/Boosting-Ensemble/decision-theoretic_generalization.pdf.
6. Breiman, L. (2001, January). *Random forest*. Berkeley, CA: Statistics Department of University of California. Retrieved from <https://www.stat.berkeley.edu/~breiman/randomforest2001.pdf>.
7. Kuznecov, I.A., & Kireev, V.S. (2016, October 11-14). Razrabotka ansamblya algoritmov klassifikacii s ispol'zovaniem entropijnogo pokazatelya kachestva dlya resheniya zadachi povedencheskogo skoringa. *Trudy XVIII Mezhdunarodnoj konferencii DAMDID/RCDL'2016 «Analitika i upravlenie dannymi v oblastyah s intensivnym ispol'zovaniem dannyh» (Yershovo, Russia), (Proceedings of XVIII International Conference «Analytics and Data Management in Data-Intensive Areas»)*, 11—42. Retrieved from <http://ceur-ws.org/Vol-1752/paper07.pdf> [in Russian].
8. Paklin, N. B., & Oreshkov, V. I. (2013). *Biznes-analitika: ot dannykh k znaniyam*. St. Petersburg, Russia: Piter [in Russian].
9. Haykin, S. (1998). *Neural Networks — A Comprehensive Foundation, Second Edition*. New Jersey: Prentice-Hall.
10. Lavrenkov Yu. N. (2014). *Issledovanie i razrabotka kombinirovanyh nejrosetevykh tekhnologij dlya povysheniya effektivnosti bezopasnoj marshrutizacii informacii v setyah svyazi: PhD Thesis*. Kaluga, Russia: MSTU named after N.E. Bauman. Retrieved from https://mpei.ru/Science/Dissertations/dissertations/Dissertations/LavrenkovYN_diss.pdf [in Russian].
11. Savina, S. S., & Ben', V. P. (2015). Objednannia modelej logit-rehesiyi yak komitetu ekspertiv dlia otsinky kredytopromozhnosti pozychal'nyka. *Nejro-nechitki tekhnologhii modeliuвання v ekonomitsi (Neuro-Fuzzy Modeling Technigues in Economics)*, 4, 154—188 [in Ukrainian].
12. Savina, S. S., & Ben', V. P. (2016). Vybir arhitektury nejromerezhi dlya rozv'yazannya zadachi klasyfikaciyi nadijnosti pozychal'nykiv-

fizychnyh osib. *Nejro-nechitki tekhnologii modeliuvannia v ekonomitsi (Neuro-Fuzzy Modeling Technigues in Economics)*, 5, 123—151 [in Ukrainian].

13. Caregorodcev, V. G. (2004). Optimizacija ekspertov boosting-kollektiva po ih krivym obuchenija. *Materialy XIII Vserossiyskogo seminar «Nejroinformatika i ee prilozhenija» (Krasnoyarsk, Russia)*, (*Proceedings of XIII All-Russian Seminar «Neuroinformatics and its applications»*), 152—157 [in Russian].

14. Schapire, R. E. (1999). Theoretical views of boosting and applications. *Proceedings of 10th International Conference «Algorithmic Learning Theory» (Tokyo, Japan)*. Retrieved from http://www-ai.cs.uni-dortmund.de/LEHRE/PG/PG445/literatur/schapire_99a.pdf.

15. Matviychuk, A. V. (2011). *Shtuchnyi intelekt v ekonomitsi: neironni merezhi, nechitka logika*. Kyiv, Ukraine: KNEU [in Ukrainian].

Стаття надійшла до редакції 04.04.2018