

О.П. Ільїна, С.Я. Скибик

АДАПТИВНА АНСАМБЛЕВА ІНТЕГРАЦІЯ РІШЕНЬ ДЛЯ ІНДИКАТОРА РЕСУРСНОЇ БЕЗПЕКИ: МЕТОДОЛОГІЯ ТА СТАТИСТИЧНА ВАЛІДАЦІЯ СТАБІЛЬНОСТІ

Забезпечення ефективної підтримки ухвалення рішень у складних розподілених організаційних системах (особливо у сфері національної безпеки та оборонного планування) вимагає розробки надійних методів класифікації, здатних оперативно діагностувати стани ресурсів та ризики стратегічним інтересам. Ефективність індикатора ресурсної безпеки (RSI), побудованого на базі методів машинного навчання, критично залежить від стабільності та достовірності інтегрованих прогнозів, особливо в умовах, характерних для предметної області: значний дисбаланс класів (де помилка пропуску негативного стану є критичною), обмежений обсяг даних, логнормальний розподіл ознак із «довгими хвостами», та наявність шумових компонентів, що знижує стабільність окремих класифікаторів. Для вирішення цього розроблено адаптивний механізм ансамблевої інтеграції (RSI), що реалізує зважене м'яке голосування моделей (NB, SVM, RF, kNN, LR) з уніфікованим калібруванням імовірностей. Центральним елементом є композитна динамічна метрика якості (KQ), яка поєднує $F1_{neg}$ (пріоритет міноритарного класу), MCC та $Kappa_{norm}$, адаптуючи їхні ваги на основі кореляції. Інтегровано коефіцієнти довіри (KDR) для корекції впливу моделей залежно від їх вразливості до властивостей даних. Валідацію алгоритму проведено на синтетичних даних, що імітують логнормальний розподіл та лагові ефекти реальних умов. Масштабний експеримент (250 прогонів, парний дизайн) підтвердив високу статистичну значущість ($p < 0.001$ за критерієм Вілкоксона) переваги RSI над найкращим окремим класифікатором (Random Forest) за всіма метриками (ΔKQ , $\Delta F1_{neg}$, $\Delta Recall_{neg}$). Розмір ефекту (d -Коена ≥ 1.41) свідчить про велику практичну цінність. Результати доводять, що адаптивна інтеграція забезпечує стабільність та надійність діагностики ризиків, що критично необхідні для безпекових застосувань.

Ключові слова: класифікація методами машинного навчання, ансамблеве навчання, адаптивна метрика якості, дисбаланс класів, м'яке голосування, статистична валідація, підтримка стратегічних рішень.

О.П. Ільїна, С.Я. Скибик

ADAPTIVE ENSEMBLE DECISION INTEGRATION FOR INDICATOR OF RESOURCE SECURITY: METHODOLOGY AND STATISTICAL VALIDATION OF STABILITY

Ensuring effective decision support in complex distributed organizational systems (especially in national security and defense planning) requires reliable classification methods capable of rapid diagnosis of resource states and risks to strategic interests. The effectiveness of a resource security indicator (RSI) built on machine learning methods critically depends on the stability and reliability of integrated predictions under conditions typical of this domain: significant class imbalance (where missing a negative state is critical), limited data volume, log-normal feature distribution with "long tails", and noise components that reduce the stability of individual classifiers. To address these challenges, an adaptive ensemble integration mechanism (RSI) was developed, implementing weighted soft voting of models (NB, SVM, RF, kNN, LR) with unified probability calibration. The central element is a composite dynamic quality metric (KQ), which combines $F1_{neg}$ (prioritizing the minority class), MCC , and $Kappa_{norm}$, adapting their weights based on correlation. Trust coefficients (KDR) are integrated to adjust the influence of models depending on their vulnerability to data properties. Algorithm validation was performed on synthetic data simulating log-normal distribution and lag effects of real-world conditions. A large-scale experiment (250 runs, paired design) confirmed high statistical significance ($p < 0.001$ by Wilcoxon test) of RSI superiority over the best single classifier (Random Forest) across all metrics (ΔKQ , $\Delta F1_{neg}$, $\Delta Recall_{neg}$). The effect size (Cohen's $d \geq 1.41$) indicates large practical value. The results demonstrate that adaptive integration ensures stability and reliability of risk diagnosis, critically necessary for security applications.

Keywords: machine-learning classification, ensemble training, adaptive quality metric, class imbalance, soft voting, statistical validation, strategic decision support.

Вступ

Для забезпечення життєздатності та захисту інтересів розподілених організаційних систем, особливо в умовах динамічного середовища воєнної сфери національної безпеки та оборонного планування, критично необхідна наявність інструментів експрес-діагностики ресурсних загроз. Розглянуте в [1] моделювання індикатора ресурсної безпеки (Resource Security Indicator — RSI) ґрунтується на сучасних методах машинного навчання (ML-класифікації), проте його ефективність на практиці стикається з низкою фундаментальних викликів, зумовлених специфікою вхідних даних. До таких викликів належать: значний дисбаланс класів на користь позитивного (нормального) стану, обмежений обсяг спостережень через специфіку доменної області, а також складність розподілів ознак. Значення останніх обмежені як доля від певного нормативу, демонструють скупчення навколо середнього або медіани, а також наявність вагомих «хвостів» внаслідок наявних стратегій постачання, що ускладнює коректне та ефективне оперування ними.

У цих умовах надійність будь-якого окремого класифікаційного механізму ставиться під сумнів, оскільки різні моделі можуть демонструвати нестабільну поведінку на різних підмножинах даних. Це робить актуальним дослідження стабільності якості класифікації та розробку інтеграційних підходів, здатних компенсувати слабкі сторони окремих методів. Динаміка змін у середовищі безпеки вимагає інструментів, які не лише надають точковий прогноз, а й адаптуються до змінних характеристик даних, мінімізуючи ризики пропуску загроз (False Negative), які можуть мати катастрофічні наслідки для стратегічного планування.

1. Аналіз літературних даних і постановка проблеми

Попередні етапи досліджень дозволили сформулювати базовий апарат побудови RSI, орієнтований на оперування за умов невизначеності, що реалізує ансамблеву інтеграцію гетерогенних

класифікаторів [1]. Однак детальний аналіз предметної області та попередніх результатів виявив необхідність орієнтувати алгоритм на доменну специфіку [2], [3] даних і пріоритетів. Це зумовило відмову від дискретного (жорсткого) голосування, де кожна модель має один рівноправний голос незалежно від ступеня її впевненості. Такий підхід втрачає критично важливу інформацію про ймовірнісний розподіл прогнозів [4]. Крім того, використання єдиної статичної метрики якості (наприклад, тільки F1 або Каппа Коена) є неадекватним для даних з високим та змінним дисбалансом. Статичні метрики часто ігнорують внутрішню кореляцію між критеріями якості, що може призводити до зміщених оцінок та «перенавчання» під мажоритарний клас, особливо коли дані містять шум або аномальні викиди.

Проблематика інтеграції рішень для RSI зосереджена навколо факторів, характерних для предметної області:

- **Різномірність ознак та складність розподілів:** Поєднання логнормальних ресурсних показників та дискретних лагових ознак ускладнює навчання стандартних моделей, вимагаючи специфічного препроцесингу.

- **Критичність гіподіагностики:** В умовах безпеки вартість пропуску негативного прецеденту (загрози) є неспівмірно вищою за вартість хибної тривоги, що вимагає жорсткої пріоритезації метрик для міноритарного класу (негативний стан).

- **Залежність моделей від властивостей даних:** Згідно з теоремою No Free Lunch [5], не існує універсального класифікатора, оптимального для всіх сценаріїв розподілу даних. Це вимагає механізму, який би динамічно враховував вразливість кожної моделі до поточних умов.

Постановка проблеми полягає у необхідності розробки та емпіричної валідації адаптивного механізму інтеграції рішень, який забезпечить стійкість, достовірність та інтерпретованість прогнозу в умовах високої невизначеності, де ціна помилки є критичною.

2. Мета і завдання дослідження

Мета дослідження: Розробити та статистично валідувати адаптивний механізм ансамблевої інтеграції рішень (RSI), здатний забезпечити стійкість і надійність оцінки в умовах, характерних для даних ресурсної безпеки (обмеженість вибірки, наявність шумів, значний дисбаланс класів).

Завдання дослідження:

1. Розробити механізм зваженого м'якого голосування (Weighted Soft Voting), що включає уніфіковане калібрування імовірностей для забезпечення математичної коректності та порівнянності виходів різнорідних моделей ансамблю.

2. Сформувати динамічну композитну метрику якості (KQ), яка автоматично адаптує ваги своїх компонентів до характеристик конкретного набору спостережень, мінімізуючи вплив мультиколінеарності критеріїв та акцентуючи увагу на виявленні загроз.

3. Статистично підтвердити перевагу адаптивного ансамблю RSI над найкращим окремим класифікатором (RF) через масштабний експеримент із використанням парного дизайну та оцінки розміру ефекту.

3. Методи і матеріали досліджень

Дослідження ґрунтується на методах ансамблевого машинного навчання, теорії математичної статистики, методах обчислювального експерименту на основі генерації синтетичних даних та розробці адаптивних метрик якості. Програмна реалізація алгоритму виконана в середовищі R v4.5.1 “Great Square Root” (реліз 13.06.2025) із використанням пакетів ``caret`` (v7.0-1, оновлено 10.12.2024), ``ranger`` (v0.17.0, оновлено 08.11.2024), ``el071`` (v1.7-16, оновлено 16.09.2024) та ``class`` (v7.3-23, оновлено 01.01.2025), що мають актуальні стабільні релізи станом на кінець 2025 року та забезпечують сучасні програмні реалізації використовуваних класифікаційних алгоритмів. Детальний опис швидкої реалізації Random Forest у пакеті ``ranger`` наведено в [6], а сучасні підходи до побудови та налаштування

моделей із використанням ``caret`` систематизовано в [7].

Архітектура гетерогенного ансамблю. RSI використовує гетерогенний набір з п'яти базових моделей: Naive Bayes (NB), Support Vector Machine (SVM), Random Forest (RF), k-Nearest Neighbors (kNN), та Logistic Regression (LR). Вибір саме цих алгоритмів обумовлений необхідністю забезпечення максимальної різноманітності помилок, що розглядається як ключова умова ефективності ансамблів [4], [8]. Наприклад, NB ефективний для незалежних ознак і швидкого навчання, SVM добре працює з високорозмірними просторами та знаходить оптимальні розділяючі гіперплощини. RF забезпечує робастність до шуму та викидів, kNN ефективний для виявлення локальних структур даних, а LR надає базову лінійну апроксимацію ймовірностей. Таке поєднання дозволяє компенсувати слабкі сторони одних моделей сильними сторонами інших, створюючи синергетичний ефект. Деталі конфігурації та гіперпараметрів цих моделей були представлені в попередній роботі [1], а узагальнений огляд ансамблевого навчання подано в [8].

Доменно-специфічні фактори.

Алгоритм RSI проєктувався як відповідь на низку викликів, характерних для ресурсних даних індикатора та контексту ухвалення рішень. Ці виклики систематизовано у вигляді множини факторів Ф1–Ф11:

Ф1. Різнорідність ознак за типами (ресурсні, лагові, службові) та шкалами вимірювання.

Ф2. Складність розподілів (обмеженість області, скупченість, правосторонні «хвости»).

Ф3. Залежність поточного стану від ресурсної передісторії (інерційність процесів, часові лаги).

Ф4. Високий дисбаланс на користь позитивного (нормального) класу.

Ф5. Критичність гіподіагностики незадовільних станів (пропуск негативних станів є набагато небезпечнішим, ніж хибна тривога).

Ф6. Наявність нефіксованих зовнішніх факторів впливу, які можуть змінювати розподіли без прямого спостереження.

Ф7. Неоднорідність умов спостережень у часі та між різними підмножинами даних.

Ф8. Епізоди оцінювання цільової змінної за відмінними експертними моделями та шкалами.

Ф9. Обмежений обсяг доступних даних у порівнянні з потенційною складністю простору ознак.

Ф10. Використання різних типів класифікаційних моделей (імовірнісних, геометричних, деревоподібних) в одному ансамблі.

Ф11. Вимога придатності результатів класифікації для практичної підтримки рішень організаційної системи.

Етапи розробки адаптивного ансамблю: Ключові етапи та дії. Запропонований алгоритм побудови та навчання індикатора становить наступну послідовність етапів та дій.

Е1. Підготовка даних

$$DSN \rightarrow LS, TS,$$

де DSN – первинний масив спостережень, LS , TS – відповідно навчальна й тестова вибірка.

Е1.1 Структуризація системи ознак.

Е1.2 Обробка пропусків.

Е1.3 Масштабування.

Е2. Навчання моделей ансамблю

$$M_i, LS \rightarrow UP_i(LS), M_i^T, KQ_i, i = 1, \dots, 6,$$

де M_i^T – тренувана на даних LS калібрована модель M_i ;

$UP_i(LS)$ – результат класифікації в точках LS від моделі M_i , поданий як калібрована імовірність позитивного класу;

KQ_i – значення трикомпонентної (критерії $F1$, MCC , $Каппа$) метрики якості класифікації.

Е2.1 Організація циклу крос-валідації на LS .

Е2.2 Параметризація процедури навчання за допомогою метрики якості.

Е2.3 Реалізація навчання M_i з використанням крос-валідації та уніфікованого калібрування результуючих

імовірностей (крос-валідаційне калібрування за методом Платта [4]).

Е3. Визначення апіорних коефіцієнтів довіри до елементів ансамблю

$$D, S, M_i \rightarrow KDR_i,$$

де D – вимоги до статистичної моделі даних, що визначають вразливості моделі M_i ;

S – ризики порушення вимог у використуваних даних;

KDR – експертна оцінка коефіцієнту довіри.

Е4. Побудова динамічної метрики якості класифікації

$$UP^{MCC}_i(LS), UP^{KAPPA}_i(LS), UP^{F1}_i(LS), KQ_i \rightarrow COR_i, W^{KQD}_i$$

де $UP^{MCC}_i(LS)$, $UP^{KAPPA}_i(LS)$, $UP^{F1}_i(LS)$ – результати класифікації, отримані аналогічно Етапу 3, але з використанням метрики якості тільки з одним із трьох елементів-критеріїв;

COR – коефіцієнт кореляції вхідних масивів;

W^{KQD} – ваговий коефіцієнт відповідного критерію в складі динамічної метрики якості KQD_i .

Е5. Прогнозна інтегрована класифікація для спостережень тестової вибірки

$$TS, \{M_i^T, W^{KQD}_i, KDR_i\}, TRP \rightarrow UP_i(TS), DP_i(TS),$$

де M_i^T – тренувані на етапі Е2 моделі;

W^{KQD} – вагові коефіцієнти критеріїв в метриках KQD моделей;

KDR – коефіцієнти довіри до моделей;

TRP – множина значень імовірності, одне з яких буде обрано як поріг для проєктування неперервних значень імовірності класу на бінарну шкалу (0, 1);

$UP_i(TS)$ – імовірнісний прогноз;

$DP_i(TS)$ – дискретизований прогноз.

Е5.1. Розрахунок ваги моделей.

Е5.2. Імовірнісний прогноз від M_i^T в точках TS .

Е5.3. Зважене усереднення прогнозу.

Е5.4. Дискретизація прогнозу.

В Табл. 1 надано обґрунтування основних положень запропонованого алгоритму.

Основні положення алгоритму

Етап	Виклики	Втілені в алгоритмі методичні рішення	Аргументація доцільності
E1.1	Ф1	Розрізнення ресурсних та лагових ознак	Різний препроцесинг для запобігання мультиколінеарності
E1.2	Ф1, Ф3	Видалення спостережень з пропуском лагів та kNN доповнення ресурсних	Збереження послідовностей
E1.3	Ф1, Ф2, Ф6, Ф10	Робастне масштабування ресурсних ознак в LR, kNN, SVM з використанням медіани та інтерквартильного розмаху (IQR) замість середнього та дисперсії (без зважування класів і видалення корельованих ознак)	Забезпечення однорідності ознак за шкалами та, що найважливіше, зменшення деструктивного впливу аномальних викидів, характерних для «важких хвостів» логнормальних розподілів.
E2.1	Ф4, Ф6, Ф7, Ф8	5 фолдів з почерговим використанням в ролі тестового	Стабільність оцінки якості
E2.2	Ф4, Ф5	Апріорна метрика якості KQ_i для моделі M_i : зважена сума адаптованих критеріїв F1 (з інверсією класів), MCC та Каппа (перенормованої)	Акцент на виявленні негативних станів, баланс аспектів якості, зіставність
E2.3	Ф10	Навчання M_i з крос-валідаційним калібруванням імовірностей за методом Платта [4]	Отримання достовірних та уніфікованих імовірностей.
E3	Ф2, Ф3, Ф6, Ф10	Визначення коефіцієнтів довіри (KDR)	Регуляція впливу моделей згідно вразливості до властивостей спостережень, ще до етапу оцінки їхньої емпіричної точності.
E4	Ф3, Ф4, Ф5	Урахування фактичної корельованості елементів KQ_i для динамічної вагової параметризації з отриманням KQD_i для M_i	Адаптивна та об'єктивна оцінка якості, що пріоритезує міноритарний клас і запобігає домінуванню колінеарних метрик
E5.1	Ф2, Ф3, Ф5, Ф10	Вагові коефіцієнти моделей M_i для наступного використання в голосуванні	Інтеграція вразливості M_i з продемонстрованою якістю
E5.2	Ф10	Прогноз індивідуальних (від M_i) імовірностей класів для тестових точок	Зіставлення M_i - прогнозів завдяки попереднім крокам
E5.3	Ф2, Ф4, Ф6, Ф7, Ф8, Ф9	Інтеграція індивідуальних прогнозів для точок тестової вибірки зваженим м'яким голосуванням [4]	Підвищення впливу переваг, запобігання переоцінці більшого класу, агрегація розмитих рішень у разі складних розподілах ознак
E5.4	Ф11	Оптимальна порогова дискретизація результатів класифікації (границя між класами встановлюється як значення ймовірності, яке максимізує KQD в точках тестової вибірки)	Створення можливості порівняння прогнозу класу з фактичними значеннями для тестової вибірки з метою наступної оцінки властивостей індикатора

Табл. 1 узагальнює основні положення алгоритму RSI та демонструє взаємозв'язок усіх його етапів із відповіддю на конкретні виклики — від підготовки даних і тренування моделей до формування та оцінки інтегрованого результату. Далі буде розглянуто низку критичних компонентів цієї схеми, які забезпечують адаптивне узгодження ансамблю з властивостями даних.

Уніфіковане калібрування рішень.

Для забезпечення сумісності ймовірностей у м'якому голосуванні використовується уніфіковане крос-валідаційне калібрування Плата. Ця процедура перетворює "сирі" вихідні дані моделей (наприклад, decision values SVM або дистанції kNN) на достовірні псевдо-ймовірності. Це особливо важливо для SVM, який стандартно повертає лише відстань до розділяючої гіперплощини, та для Random Forest, який часто схильний до надмірної впевненості на краях діапазону (близько 0 та 1). Калібрування дозволяє привести виходи всіх моделей до єдиної ймовірнісної шкали, що є необхідною умовою для коректного математичного усереднення; детальний аналіз підходів до калібрування наведено в [9], [10].

Визначення коефіцієнту довіри KDR. Вплив моделей в ансамблі додатково коригується коефіцієнтом довіри KDR_i , який відображає ступінь порушення припущень даних для кожної моделі. Розрахунок KDR_i дозволяє інтегрувати апіорні експертні знання про вразливість моделей згідно з формулою:

$$KDR_i = \frac{1}{4} \sum_j EKD[i, j] \times IV[i, j]$$

де EKD — матриця експертних оцінок вразливості моделей до порушення властивостей даних, а

IV — індикатори фактичного порушення цих умов у поточному наборі спостережень.

Для кожного класифікатора MC_i з пулу $PC = \{NB, SVM, RF, kNN, LR\}$ на основі аналізу широкого спектру класичних публікацій, зокрема [11], експертно оцінюються критичність чотирьох вимог до даних:

1. Форма розподілу ознак (feature distribution);
2. Незалежність ознак (feature independence);
3. Баланс спостережень за класами (observation data balance);
4. Шкала вимірювання (measurement scale).

Кожна вимога j для моделі i оцінюється на тривірневій шкалі $ekd_{ij} \in \{0; 0.5; 1\}$, де 0 означає відсутність суттєвих вимог моделі до відповідної властивості даних, 0.5 — помірну чутливість (порушення припущення допустимі за наявності стандартних компенсуючих заходів), 1 — критичну чутливість. Обґрунтування кожного значення ekd_{ij} проводилося за чотирма аспектами $A_{a,j}$ (із $a = 1, \dots, 4$):

- ($a = 1$) - наявність первинного методичного обмеження;
- ($a = 2$) - доступні та фактично застосовані в RSI заходи робастності;
- ($a = 3$) - програмні засоби (R-паке-ти/функції), використані для реалізації цих заходів;
- ($a = 4$) - потенційні ризики для якості класифікації у разі порушення вимоги.

Підсумкова матриця EKD , наведена у Табл. 2, фіксує лише *теоретико-методичну вразливість моделей*; у формулі для KDR_i вона поєднується з масивом індикаторів IV_{ij} , які залежать від конкретних спостережень.

Експертні коефіцієнти критичності ekd_{ij} та аргументація для моделей ансамблю RSI

	Розподіл ознак A_{a1}	Незалежність ознак A_{a2}	Збалансованість спостережень A_{a3}	Шкала значень A_{a4}
Байєсівський класифікатор				
ekd_{i1}	0	0	0.5	0
$a = 1$	Не накладаються жорсткі параметричні вимоги до форми розподілу ознак	Передбачається умовна незалежність ознак	Дисбаланс класів зміщує апостеріорні ймовірності до мажоритарного класу	Окремих вимог немає за умови коректного препроцесингу
$a = 2$	Непараметрична оцінка щільності (KDE) із параметрами $usekernel = TRUE$, $adjust = 1$	Ознаки структуровано за типами (ресурсні, лагові), що обмежує складні залежності	Баланс класів частково стабілізується завдяки стратифікованому поділу train/test	Працює з уже узгоджено препроцесованими ознаками
$a = 3$	<code>caret::train (method="nb"); e1071::naiveBayes</code>	<code>caret, e1071</code>	<code>caret::confusionMatrix; mltools</code>	<code>caret, e1071</code>
$a = 4$	Мінімальні	Сильні приховані залежності між ознаками можуть порушувати припущення умовної незалежності	Ризик пропуску негативних станів	Непоследовне кодування типів ознак може викривляти інтерпретацію щільностей
Метод опорних векторів SVM				
ekd_{i2}	0	0	0.5	0.5
$a = 1$	Немає вимоги нормальності; проблеми переважно за умови дуже складних кордонах та шумі	Загалом стійка до помірної мультиколінеарності	За суттєвого дисбалансу гіперплощина може «зміщуватися» на користь мажоритарного класу	Дуже чутлива до масштабу ознак
$a = 2$	Застосовано радіальне ядро з підбором параметрів регуляризації та ширини ядра	Ознаки структуровано за типами (ресурсні, лагові), що обмежує складні залежності	Баланс класів частково стабілізується завдяки стратифікованому поділу train/test	Робастне масштабування числових ознак (центрування за медіаною та масштабування за IQR)
$a = 3$	<code>caret::train (method="svmRadial"); kernlab::ksvm; stats::glm</code>	<code>caret, kernlab, stats::glm</code>	<code>caret, kernlab</code>	<code>caret, kernlab</code>

	Розподіл ознак A_{a1}	Незалежність ознак A_{a2}	Збалансованість спостережень A_{a3}	Шкала значень A_{a4}
$a = 4$	Невдалий вибір параметрів ядра може призводити до пере- або недонавчання на складних розподілах	Мультиколінеарність здатна погіршувати стабільність опорних векторів і ускладнювати інтерпретацію моделі	За значного дисбалансу можливе зміщення розділяючої гіперплощини на користь мажоритарного класу	Некоректне масштабування ознак здатне суттєво викривлювати відстані в ознаковому просторі
Випадковий ліс				
ekd_{i3}	0	0	0.5	0
$a = 1$	Відсутні	Відсутні	Схильність до мажоритарного класу при дисбалансі	Відсутні
$a = 2$	Непотрібні	Ознаки групуються за типами (ресурсні, лагові)	Баланс класів частково стабілізується завдяки стратифікованому поділу train/test	Не застосовується
$a = 3$	<code>caret::train (method= "rf"); randomForest; ranger</code>			
$a = 4$	Відсутні	Відсутні	За сильного дисбалансу існує ризик систематичного недопредставлення міноритарного класу	Відсутні
Метод найближчих сусідів kNN				
ekd_{i4}	0	0.5	0.5	0.5
$a = 1$	Не висуваються припущення щодо форми розподілів, але є чутливість до шуму	Сильні кореляції між ознаками спотворюють евклідові відстані	У разі дисбалансу найближчі сусіди часто належать мажоритарному класу	Відстані сильно залежать від масштабу ознак; «довгі» шкали домінують у метриці
$a = 2$	Не застосовуються; робастність досягається за рахунок вибору відстаневої метрики та масштабування	Ознаки розділено за типами (ресурсні, лагові)	Баланс класів частково стабілізується завдяки стратифікованому поділу train/test	Для числових ознак застосовано робастне масштабування (медіана/IQR), що зменшує домінування окремих ознак у відстанях
$a = 3$	<code>caret::train (method= "knn"); class::knn; kknn::train.kknn</code>			

	Розподіл ознак A_{a1}	Незалежність ознак A_{a2}	Збалансованість спостережень A_{a3}	Шкала значень A_{a4}
$a = 4$	Наявність шуму та розрідженості може робити сусідства нестабільними	Сильні кореляції між ознаками спотворюють геометрію простору	Виражений дисбаланс класів може призводити до домінування мажоритарного класу серед найближчих сусідів	Відсутність коректного масштабування робить результат залежним від вибору одиниць вимірювання
Логістична регресія				
ekd_{15}	0	0.5	0.5	0.5
$a = 1$	Не вимагається, проте викиди та асиметрія можуть впливати на стабільність оцінок	Мультиколінеарність збільшує дисперсію оцінених коефіцієнтів та ускладнює інтерпретацію	У разі дисбалансу класів максимізація правдоподібності тяжіє до мажоритарного класу	Чутливість до некоректного кодування категоріальних змінних та різких відмінностей масштабів
$a = 2$	Для зменшення впливу викидів використано робастне масштабування числових ознак	Ознаки розділено за типами (ресурсні, лагові)	Баланс класів частково стабілізується завдяки стратифікованому поділу train/test	Для числових ознак використано робастне масштабування (медіана/IQR), що підвищує стабільність оцінених коефіцієнтів
$a = 3$	<code>caret::train (method= "glm", family =binomial); stats::glm</code>			
$a = 4$	Сильні відхилення від «регулярних» розподілів (викиди, важкі «хвости») можуть змішувати оцінки параметрів	Виражена мультиколінеарність здатна значно збільшувати дисперсію оцінених коефіцієнтів	У разі значного дисбалансу логістична регресія може повертати добре калібровані, але зміщені ймовірності з низькою чутливістю до міноритарного класу	Некоректне кодування категоріальних змінних і великі відмінності масштабів між ознаками можуть суттєво погіршувати якість класифікації

Адаптивна метрика якості (KQ).

Якість класифікації оцінюється за композитним індикатором KQ , розробленим у цьому дослідженні спеціально для обробки незбалансованих даних; детальний огляд проблеми навчання на незбалансованих вибірках і відповідних метрик подано, зокрема, в [12], [13], [14]. Метрика включає три компоненти:

- $F1_{neg}$ (базова вага $w_{f1} = 0.5$) — гармонічне середнє Precision та Recall для негативного класу. Висока вага цього

компонента є свідомим вибором для пріоритезації виявлення критичних загроз, навіть якщо це призводить до незначного зниження загальної точності [2];

- MCC (базова вага $w_{mcc} = 0.35$) — коефіцієнт кореляції Метьюса. Це робастна оцінка загальної якості, яка враховує всі чотири елементи матриці плутанини (TP, TN, FP, FN) і залишається адекватною навіть за сильного дисбалансу класів;

- $Kappa_{norm}$ (базова вага $w_{kappa} = 0.15$) — нормалізована каппа Коена, що

відображає перевагу моделі над випадковим вгадуванням. Метрика є динамічною, оскільки її вагові коефіцієнти автоматично коригуються на основі кореляції r між MCC та $Kappa_{norm}$ (фактор надлишковості). Якщо ці метрики сильно корелюють на поточному наборі даних, їхні ваги пропорційно зменшуються, щоб уникнути дублювання інформації та перекосу оцінки. Це забезпечує адаптованість KQ до фактичних характеристик навчальних даних і робить її більш стійкою до змін у розподілах.

Зважене м'яке голосування. Фінальне інтегроване рішення отримується через процедуру зваженого м'якого голосування [4], [8]:

$$\hat{p}_i = \sum_{m=1}^N W_m \cdot P_m(x_i),$$

де $P_m(x_i)$ — калібрована імовірність позитивного класу від моделі m , а

вага W_m залежить від якості моделі KQ_m та її коефіцієнта довіри KDR_m : $W_m \propto KQ_m \cdot (1 - KDR_m)$.

Обґрунтування переваги м'якого голосування полягає у використанні всієї інформації про ступінь впевненості моделі. На відміну від жорсткого голосування, де "невпевнене" рішення моделі (наприклад, ймовірність 0.51) має таку ж вагу, як і "впевнене" (0.99), м'яке голосування дозволяє точніше агрегувати прогнози, надаючи більшу вагу більш впевненим моделям. Це призводить до формування гладкішої розділяючої поверхні та значно якісніших інтегрованих імовірностей. Оптимальний поріг дискретизації встановлюється не фіксовано (0.5), а адаптивно — як значення, що максимізує фінальну метрику KQ на тестовій вибірці, що дозволяє гнучко балансувати точність і повноту залежно від поточних умов задачі.

Генерація даних для експерименту. Для дослідження стабільності використано синтетично згенеровані дані, що імітують складні виклики предметної області: логнормальний розподіл ознак, дисбаланс (20% негативного класу) та інерційність процесів (лагові

ознаки). Значення мітки імітують експертні оцінки цільової благополучності, які надаються наприкінці чергового періоду аудиту та розповсюджуються на ресурсні спостереження по всій його довжині. Введення лагових ознак паралельно з ознаками рівня забезпеченості за видами ресурсу обслуговує потребу в урахуванні впливів на благополучність, спричинених ресурсними даними попередніх періодів. Використання синтетичних даних є необхідним методологічним кроком, оскільки дозволяє точно контролювати параметри розподілів та рівень шуму, що неможливо на обмежених реальних історичних вибірках. Експеримент проводився на 5 серіях датасетів з обсягами вибірок: 3000, 1500, 840, 364 та 140 спостережень. Логіка генератора ґрунтується на «м'якому скоринговому відборі» (soft scoring selection). Цей метод враховує поточний стан системи (L), попередній стан (L_{prev}) та інерційні параметри ($FL1$ — довжина серії станів) для відбору значень з батьківських логнормальних розподілів. Скоринг-функція S інтегрує вплив «хвоста» розподілу та довжини серії періодів із стабільним значенням мітки, що дозволяє генерувати реалістичні часові ряди з залежностями, імітуючи складні процеси деградації ресурсів, а не просто випадковий шум.

Дизайн статистичного експерименту.

Експеримент включав 250 незалежних прогонів на 5 різних синтетичних датасетах. Було застосовано парний дизайн (Paired Design) для порівняння ансамблю RSI проти Random Forest (RF) [15]. Вибір RF як базового еталону є обґрунтованим, оскільки попередні дослідження показали, що RF є найкращим окремим класифікатором серед компонентів ансамблю. Таким чином, це найбільш суворий тест, що дозволяє оцінити саме додаткову цінність механізму інтеграції, а не просто перевагу над слабкими моделями. Парний дизайн нівелює вплив випадкових факторів (випадкове початкове значення, розбиття на навчальну/тестову вибірки), фокусуючись виключно на різниці в ефективності методів на одних і тих самих даних. Для оцінки переваги використовувався аналіз

дельт ($\Delta = RSI - RF$) та такі статистичні критерії:

- *Критерій знакових рангів Вілкоксона* — непараметричний тест для перевірки гіпотези про зсув медіани різниць;
- *Парний t-критерій Стьюдента* — для перевірки середнього значення різниць;
- *Розмір ефекту d-Коена* — для оцінки практичної значущості отриманої різниці.

4. Результати досліджень

Огляд загальної ефективності ансамблю. Проведені дослідження на 250

прогонах показали, що ансамбль RSI перевершив якість найкращої окремої моделі (Random Forest) у всіх 5 серіях експерименту (з обсягами від 3000 до 140 спостережень), незалежно від параметрів датасету (розміру вибірки та рівня шуму). Це свідчить про універсальність запропонованого підходу.

Статистична значущість переваги RSI. Середні абсолютні значення метрик якості, отримані під час експерименту, наведені у Табл. 3. Результати аналізу дельт ($\Delta = RSI - RF$) за загальною вибіркою (250 прогонів) підтвердили статистичну значущість переваги RSI, що детально представлено у Табл. 4.

Таблиця 3

Середні абсолютні значення метрик (250 прогонів)

Метрика	Ансамбль RSI (Середнє)	Random Forest (Середнє)	Різниця Δ (RSI - RF)	Медіана Δ (RSI - RF)
KQ	0.6260	0.5578	+0.0682	+0.0499
$F1_{neg}$	0.6432	0.5609	+0.0823	+0.0647
$Recall_{neg}$	0.7108	0.4646	+0.2462	+0.2143

Таблиця 4

Статистична значущість переваги RSI над RF (250 прогонів)

Метрика	Тест Вілкоксона, p	Тест Вілкоксона, 95% ДІ	t-критерій Стьюдента, p	t-критерій Стьюдента, 95% ДІ	Розмір ефекту d - Коена
ΔKQ	2.06×10^{-38} (< 0.001)	[0.0528, ∞]	1.20×10^{-39} (< 0.001)	[0.0610, ∞]	1.41
$\Delta F1_{neg}$	3.40×10^{-41} (< 0.001)	[0.0653, ∞]	4.63×10^{-47} (< 0.001)	[0.0748, ∞]	1.60
$\Delta Recall_{neg}$	1.28×10^{-41} (< 0.001)	[0.2222, ∞]	4.49×10^{-70} (< 0.001)	[0.2299, ∞]	2.23

Ключові висновки з отриманих результатів:

- **Статистична значущість:** Усі p -значення для обох статистичних тестів

(Вілкоксона та Стьюдента) значно нижчі за рівень значущості $\alpha = 0.05$ (фактично $p < 10^{-38}$). Це доводить статистично значущу перевагу ансамблю RSI над RF. Імовірність

того, що цей результат є випадковим збігом, фактично нульова.

- **Практична значущість:** Розмір ефекту d-Коена для всіх ключових метрик є великим ($d \geq 1.41$). Отримані значення $d > 1.4$ свідчать про те, що різниця між RSI та RF є не тільки статистично фіксованою, а також обґрунтованою та помітною на практиці. Це означає, що покращення якості класифікації є суттєвим і стабільним, а розподіли результатів двох методів мають мале перекриття. Особливо важливим є значне покращення $\Delta Recall_{neg}$ (+24.6%), що вказує на здатність ансамблю значно краще виявляти приховані загрози.

5. Обговорення результатів

Отримані результати підтверджують ключову гіпотезу дослідження: адаптивна інтеграція рішень є ефективним способом підвищення надійності та стабільності класифікації в умовах, які є традиційно складними для окремих ML-моделей.

Додаткова цінність ансамблювання. Той факт, що ансамбль RSI статистично значуще перевершив свій найкращий окремий компонент (Random Forest), доводить, що механізм інтеграції збільшує додаткову цінність, яка виправдовує підвищену обчислювальну складність алгоритму. Ця перевага виникає завдяки синергії двох ключових факторів:

- **роль гетерогенності та динамічної адаптації:** Хоча RF є найкращим у середньому, у певних підмножинах даних (наприклад, на краях розподілів або при специфічних комбінаціях шуму) якісний результат можуть демонструвати інші моделі (SVM, NB). Адаптивна інтеграція (через KDR-зважування та динамічний KQ) дозволяє RSI динамічно «нахилятися» до тимчасово найкращої моделі, уникнувши пастки локального оптимуму одного методу. Це забезпечує робастність системи.

- **робастність м'якого голосування:** Зважене м'яке голосування, підкріплене уніфікованим калібруванням, забезпечило стабільну перевагу над RF в умовах обмежених, зашумлених та значно незбалансованих наборів даних. Цей підхід дозволяє врахувати нюанси розподілу

ймовірностей, де «впевнені» прогнози (наприклад, з імовірністю 0.9) вносять суттєво більший вклад у фінальне рішення, ніж «граничні» (0.51). Це підтверджує, що агрегація ступеня впевненості моделей формує більш стійкий вирішальний простір і є значно інформативнішою, ніж проста мажоритарна агрегація бінарних рішень, яка втрачає цю критично важливу мета-інформацію.

Це підтверджує, що розроблений ансамбль не просто усереднює результати, а створює адаптивний механізм, який мінімізує вплив нестабільності окремих моделей і забезпечує вищу надійність індикатора в цілому.

6. Висновки

Отримані результати підтверджують, що запропонований алгоритм адаптивної ансамблевої інтеграції рішень є перспективним підходом для побудови індикатора ресурсної безпеки інтересів розподіленої організаційної системи:

1. Розроблена методологія успішно поєднує зважене м'яке голосування, уніфіковане калібрування імовірностей, адаптивну композитну метрику якості (KQ) та експертно задані коефіцієнти відповідності даним (EKD, KDR), що дозволяє ефективно працювати зі складними незбалансованими даними та враховувати теоретичну вразливість моделей до порушення статистичних припущень.

2. Проведений статистичний експеримент показав наявність доменно репрезентативних областей даних зі статистично значущою перевагою ансамблю RSI над еталонним класифікатором RF ($p < 0.001$), підтвердивши гіпотезу про ефективність адаптивного підходу.

3. Розмір ефекту (d-Коена) підтвердив практичну значущість переваги ($d > 1.4$), довівши, що механізм інтеграції створює реальну додаткову цінність.

4. Забезпечення стабільності та надійності класифікації робить RSI ефективним інструментом для застосування у сфері національної безпеки та оборонного планування, де ціна помилки в бік

нехтування ресурсною неблагополучністю є критично високою.

5. Реалізація RSI базується на сучасному стеку пакетів R (``caret``, ``ranger``, ``e1071``, ``class``) з актуальними версіями станом на кінець 2025 року, що додатково підкріплює відтворюваність результатів і відповідність застосованих програмних засобів сучасним практикам побудови класифікаційних моделей.

Напрямки подальших досліджень.

Подальші дослідження будуть зосереджені на системному аналізі стабільності досягнутої якості класифікації (метрика KQ) відносно ключових детермінант проблемних даних. Планується використання методології Latin Hypercube Sampling (LHC) для ефективного покриття багатовимірного простору факторних викликів (рівня шуму, ступеня дисбалансу, об'єму даних) та отримання кількісних оцінок меж стабільності алгоритму в ширшому діапазоні умов експлуатації, як показано в [16].

Література

1. Skybyk S., Doroshenko A., Ilyina O., Sinitsyn I. Machine-Learning-Based Model for Indicators of the Resource-Based Security of Interests in High-Level Organizational Systems // *Proceedings of the 9th International Scientific and Practical Conference Applied Information Systems and Technologies in the Digital Society (AISTDS 2025)*. CEUR Workshop Proceedings. 2025. Vol. 4133. P. 153–169. URL: https://ceur-ws.org/Vol-4133/S_12_Doroshenko.pdf.
2. Kress M. *Operational Logistics: The Art and Science of Sustaining Military Operations*. 2nd ed. Switzerland: Springer International Publishing, 2016. 313 p.
3. Сініцин І.П., Шевченко В.Л., Дорошенко А.Ю., Федоренко Р.М. Моделі та програмні системи управління оборонними ресурсами: монографія. Київ: ІПС НАНУ, 2024. 268 с.
4. Large J., Lines J., Bagnall A. A probabilistic classifier ensemble weighting scheme based on exponentially weighting the probability estimates // *Data Mining and Knowledge Discovery*. 2019. DOI: 10.1007/s10618-019-00638-y.
5. Adam S. P., Alexandropoulos S.-A. N., Pardalos P. M., Vrahatis M. N. No Free Lunch Theorem: A review // In: Demetriou I. C., Pardalos P. M. (eds.) *Approximation and Optimization*. Springer Optimization and Its Applications, vol. 145. Switzerland: Springer, 2019. P. 57–82. DOI: 10.1007/978-3-030-12767-1_5.
6. Wright M. N., Ziegler A. ranger: A Fast Implementation of Random Forests for High Dimensional Data in C++ and R // *Journal of Statistical Software*. 2017. Vol. 77, No. 1.
7. Kuhn M., Johnson K. *Feature Engineering and Selection: A Practical Approach for Predictive Models*. Boca Raton, FL: CRC Press, 2019. 600 p.
8. Kuncheva L. I. *Combining Pattern Classifiers: Methods and Algorithms*. Wiley, 2004. 312 p.
9. Kull M., Silva F. A., Flach P. Beyond Sigmoids: How to Obtain Well-Calibrated Probabilities from Binary Classifiers // *Machine Learning*. 2017. Vol. 106, No. 3. P. 437–451.
10. Naeini M.P., Cooper G.F., Hauskrecht M. Obtaining Well Calibrated Probabilities Using Bayesian Binning // *Proceedings of the AAAI Conference on Artificial Intelligence*. 2015.
11. Aggarwal C. C. *Data Mining: The Textbook*. Springer, 2015. 734 p.
12. He H., Ma Y. *Imbalanced Learning: Foundations, Algorithms, and Applications*. Wiley, 2013. DOI: 10.1002/9781118646106.
13. Branco P., Torgo L., Ribeiro R.P. A Survey of Predictive Modeling Under Imbalanced Distributions // *ACM Computing Surveys*. 2015. Vol. 49, No. 2.
14. Krawczyk B. Learning from Imbalanced Data: Open Challenges and Future Directions // *Progress in Artificial Intelligence*. 2016. Vol. 5, No. 4. P. 221–232.
15. Demšar J. Statistical Comparisons of Classifiers over Multiple Data Sets // *Journal of Machine Learning Research*. 2006. Vol. 7. P. 1–30.
16. Iooss B., Lemaître P. A Review on Global Sensitivity Analysis Methods // In: Dellino G., Meloni C. (eds.) *Uncertainty Management in Simulation-Optimization of Complex Systems*. Boston, MA: Springer, 2015. P. 101–122.

References

1. Skybyk S., Doroshenko A., Ilyina O., Sinitsyn I. Machine-Learning-Based Model for Indicators of the Resource-Based Security of Interests in High-Level Organizational

- Systems // Proceedings of the 9th International Scientific and Practical Conference *Applied Information Systems and Technologies in the Digital Society* (AISTDS 2025). CEUR Workshop Proceedings. 2025. Vol. 4133. P. 153–169. URL: https://ceur-ws.org/Vol-4133/S_12_Doroshenko.pdf.
2. Kress M. Operational Logistics: The Art and Science of Sustaining Military Operations. 2nd ed. Switzerland: Springer International Publishing, 2016. 313 p.
 3. Sinitsyn I. P., Shevchenko V. L., Doroshenko A. Yu., Fedorenko R. M. Models and software systems for defence resource management: monograph. Kyiv: IPS of NASU, 2024. 268 p.
 4. Large J., Lines J., Bagnall A. A probabilistic classifier ensemble weighting scheme based on exponentially weighting the probability estimates // *Data Mining and Knowledge Discovery*. 2019. DOI: 10.1007/s10618-019-00638-y.
 5. Adam S. P., Alexandropoulos S.-A. N., Pardalos P. M., Vrahatis M. N. No Free Lunch Theorem: A review // In: Demetriou I. C., Pardalos P. M. (eds.) *Approximation and Optimization*. Springer Optimization and Its Applications, vol. 145. Switzerland: Springer, 2019. P. 57–82. DOI: 10.1007/978-3-030-12767-1_5.
 6. Wright M. N., Ziegler A. ranger: A Fast Implementation of Random Forests for High Dimensional Data in C++ and R // *Journal of Statistical Software*. 2017. Vol. 77, No. 1.
 7. Kuhn M., Johnson K. Feature Engineering and Selection: A Practical Approach for Predictive Models. Boca Raton, FL: CRC Press, 2019. 600 p.
 8. Kuncheva L. I. Combining Pattern Classifiers: Methods and Algorithms. Wiley, 2004. 312 p.
 9. Kull M., Silva F. A., Flach P. Beyond Sigmoids: How to Obtain Well-Calibrated Probabilities from Binary Classifiers // *Machine Learning*. 2017. Vol. 106, No. 3. P. 437–451.
 10. Naeini M.P., Cooper G.F., Hauskrecht M. Obtaining Well Calibrated Probabilities Using Bayesian Binning // *Proceedings of the AAAI Conference on Artificial Intelligence*. 2015.
 11. Aggarwal C. C. Data Mining: The Textbook. Springer, 2015. 734 p.
 12. He H., Ma Y. Imbalanced Learning: Foundations, Algorithms, and Applications. Wiley, 2013. DOI: 10.1002/9781118646106.
 13. Branco P., Torgo L., Ribeiro R.P. A Survey of Predictive Modeling Under Imbalanced Distributions // *ACM Computing Surveys*. 2015. Vol. 49, No. 2.
 14. Krawczyk B. Learning from Imbalanced Data: Open Challenges and Future Directions // *Progress in Artificial Intelligence*. 2016. Vol. 5, No. 4. P. 221–232.
 15. Demšar J. Statistical Comparisons of Classifiers over Multiple Data Sets // *Journal of Machine Learning Research*. 2006. Vol. 7. P. 1–30.
 16. Iooss B., Lemaître P. A Review on Global Sensitivity Analysis Methods // In: Dellino G., Meloni C. (eds.) *Uncertainty Management in Simulation-Optimization of Complex Systems*. Boston, MA: Springer, 2015. P. 101–122.

Одержано: 30.11.2025

Внутрішня рецензія отримана: 07.12.2025

Зовнішня рецензія отримана: 11.12.2025

Про авторів:

Скибик Сергій Ярославович,
аспірант
<https://orcid.org/0009-0008-4336-680X>

Ільїна Олена Павлівна,
кандидат фізико-математичних наук,
старший науковий співробітник,
провідний науковий співробітник
<https://orcid.org/0000-0002-4073-9649>

Місце роботи авторів:

Інститут програмних систем
Національної академії наук України,
03187, м. Київ, проспект
Академіка Глушкова, 40.
e-mail: ilyina.elena1@ukr.net,
sskybyk@gmail.com