

7.4 Методологічні засади інтелектуальної обробки даних в інтелектуальних системах підтримки прийняття рішень

Вступ

Зростання обсягів інформації, що циркулює в різноманітних системах збору, обробки та передачі інформації призводить до значного використання обчислювальних ресурсів апаратних засобів. Збройні сили технічно розвинених країн мають інтегровані архітектури прийняття рішень, що базується на:

- штучному інтелекті та нанотехнологіях;
- ефективній обробці великих масивів інформації;
- багатофункціональних процесорах зі здатністю підтримки прийняття рішень у реальному масштабі часу;
- технологіях стиснення даних для підвищення швидкості їх обробки.

При цьому використання інформаційних систем з елементами штучного інтелекту дозволить підвищити ефективність операцій (бойових дій), вплине на доктрину, організацію та способи застосування угруповань військ (сил).

Разом з тим підвищення динамічності операцій, зростання кількості різноманітних сенсорів та необхідність інтеграції їх у єдиний інформаційний простір створює ряд проблем:

- реалізовані алгоритми встановлення кореляцій між подіями недостатньо повно враховують надійність джерел розвідувальних відомостей і достовірність інформації в динаміці бойових дій;
- форми представлення інформації ускладнюють її передачу по каналам зв'язку;
- обмежені обчислювальні потужності апаратних засобів;
- обмежена пропускну здатність каналів передачі даних;
- радіоелектронне подавлення каналів КХ та УКХ радіозв'язку та кібернетичний вплив на інформаційні системи.

– перехід до принципу оцінки об’єктів моніторингу “усе впливає на все й відразу”, яке охоплює сукупні мережні та обчислювальні ресурси всіх видів збройних сил.

Саме тому, необхідно проводити розробку алгоритмів (методів та методик) які здатні за обмежений час та з високим ступенем достовірності провести оцінку стану об’єкту моніторингу від різнотипних джерел розвідувальних відомостей.

Враховуючи зазначене, актуальним науковим завданням є розробка методу оцінки стану об’єкту моніторингу в інформаційних системах спеціального призначення.

Аналіз літературних даних та постановка проблеми

В роботі [352] представлений алгоритм когнітивного моделювання. Визначено основні переваги когнітивних інструментів. До недоліків зазначеного підходу слід віднести відсутність врахування типу невизначеності про стан об’єкту аналізу.

В роботі [353] розкрито суть когнітивного моделювання та сценарного планування. Запропонована система взаємодоповнюючих принципів побудови і реалізації сценаріїв, виділені різні підходи до побудови сценаріїв, описана процедура моделювання сценаріїв на основі нечітких когнітивних карт. Запропонований авторами підхід не дозволяє врахувати тип невизначеності про стан об’єкту аналізу та не враховує зашумленість початкових даних.

В роботі [354] проведений аналіз основних підходів до когнітивного моделювання. Когнітивний аналіз і дозволяє: дослідити проблеми з нечіткими чинниками і взаємозв’язками; враховувати зміни зовнішнього середовища та використовувати об’єктивно сформовані тенденції розвитку ситуації в своїх інтересах. Разом з тим, в зазначеній роботі не дослідженим залишається питання опису складних та динамічних процесів.

В роботі [355] представлено метод аналізу великих масивів даних. Зазначений метод орієнтований на пошук скритої інформації в великих масивах даних. Метод включає операції генерування аналітичних базових ліній, зменшення змінних, виявлення розріджених ознак та наведення правил. До

недоліків зазначеного методу належить неможливість врахування різних стратегій оцінювання рішень, відсутність врахування типу невизначеності вхідних даних.

В роботі [356] наведений механізм трансформації інформаційних моделей об'єктів будівництва до їх еквівалентних структурних моделей. Цей механізм призначений для автоматизації необхідних операцій з перетворення, модифікації та доповнення під час такого обміну інформацією. До недоліків зазначеного підходу слід віднести неможливість оцінити адекватність та достовірність процесу трансформації інформації, а також провести відповідну корекцію отриманих моделей.

В роботі [357] проведено розробку аналітичної web-платформи для дослідження географічного та часового розподілу інцидентів. Web-платформу, містить декілька інформаційних панелей зі статистично значущими результатами за територіями. До недоліків зазначеної аналітичної платформи належить неможливість оцінити адекватність та достовірність процесу трансформації інформації, а також висока обчислювальна складність. Також до недоліків зазначеного дослідження слід віднести не однонаправленість пошуку рішення.

В роботі [358] проведено розробку методу нечіткого ієрархічного оцінювання якості обслуговування бібліотек. Зазначений метод дозволяє провести оцінювання якості бібліотек за множиною вхідних параметрів. До недоліків зазначеного методу слід віднести неможливість оцінити адекватність та достовірність оцінки та відповідно визначити похибку оцінювання.

В роботі [359] проведено аналіз 30 алгоритмів обробки великих масивів даних. Показано їх переваги та недоліки. Встановлено, що аналіз великих масивів даних повинен проводитися пошарово, відбуватися в режимі реального часу та мати можливість до самонавчання. До недоліків зазначених методів слід віднести їх велику обчислювальну складність та неможливість провести перевірку адекватності отриманих оцінок.

В роботі [360] представлено підхід з оцінки вхідних даних для систем

підтримки та прийняття рішень. Сутність запропонованого підходу полягає в кластеризації базового набору вхідних даних, їх аналізу, після чого на підставі аналізу відбувається навчання системи. Недоліками зазначеного підходу є поступове накопичення помилки оцінювання та навчання в зв'язку з відсутністю можливості оцінки адекватності прийнятих рішень.

В роботі [361] представлено підхід щодо обробки даних з різних джерел інформації. Зазначений підхід дозволяє проводити обробку даних з різних джерел. До недоліків зазначеного підходу слід віднести низьку точність отриманої оцінки та неможливість здійснити перевірку достовірності отриманої оцінки.

В роботі [362] проведений порівняльний аналіз існуючих технологій підтримки прийняття рішень, а саме: метод аналізу ієрархій, нейронні мережі, теорія нечітких множин, генетичні алгоритми і нейро-нечітке моделювання. Вказані переваги і недоліки даних підходів. Визначено сфери їх застосування. Показано, що метод аналізу ієрархій добре працює за умови повної початкової інформації, але в силу необхідності порівняння експертами альтернатив і вибору критеріїв оцінки має високу частку суб'єктивізму. Для задач прогнозування в умовах ризику і невизначеності обґрунтованим є використання теорії нечітких множин і нейронних мереж.

В роботі [363] розроблено метод структурно-цільового аналізу розвитку слабоструктурованих систем: підхід до дослідження конфліктних ситуацій, що викликані протиріччями в інтересах суб'єктів, які впливають на розвиток досліджуваної системи і методи вирішення слабоструктурованих проблем на підставі формування сценаріїв розвитку ситуації. При цьому проблема визначається як невідповідність існуючого стану системи необхідному, яке задано суб'єктом управління. Разом з тим, до недоліків запропонованого методу слід віднести проблему локального оптимуму та неможливість проведення паралельного пошуку.

В роботі [364] представлено когнітивний підхід до імітаційного моделювання складних систем. Показано переваги зазначеного підходу, який

дозволяє описати ієрархічні складі системи. До недоліків запропоновано підходу слід віднести відсутність врахування обчислювальних ресурсів системи.

Проведення аналізу праць [352–364] показав що спільними недоліками вищезазначених досліджень є:

- відсутність можливості формування ієрархічної системи показників;
- відсутність врахування обчислювальних ресурсів системи;
- відсутність механізмів корегування системи показників в ході оцінювання;
- відсутність механізмів глибокого навчання баз знань;
- відсутність врахування обчислювальних ресурсів, доступних в системі.

З цією метою пропонується провести розробку методу оцінки в інформаційних системах спеціального призначення

Метою дослідження є розробка методу оцінки стану об'єкту моніторингу в інформаційних системах спеціального призначення, який б дозволив проводити аналіз стану об'єктів з заданою достовірністю при ресурсних обмеженнях.

Для досягнення мети були поставлені такі завдання:

- сформулювати концепцію представлення методу оцінки стану об'єкту моніторингу в інформаційних системах спеціального призначення;
- визначити алгоритм реалізації методу;
- навести приклад застосування запропонованого методу при аналізі оперативної обстановки угруповання військ (сил).

Результати дослідження з розробки методу оцінки стану об'єкту моніторингу в інформаційних системах спеціального призначення

Концепція представлення методу оцінки стану об'єкту моніторингу в інформаційних системах спеціального призначення

Систему управління процесом аналізу стану об'єктів можна представити у вигляді знакового орієнтованого графа. Загалом завдання визначення стану об'єкту моніторингу зводиться до розрахунків відповідно до формули:

$$A_i(k+1) = f \left(\left(A_i(k) + \sum_{j \neq i, j=1}^N A_j(k) W_{ij} \right) \times t_{ij} \right) \times \zeta_{ij}, \quad (1)$$

де $A_i(k+1)$ – новий стан вершини графа, $A_i(k)$ – попередній стан графа, W_{ij} – матриця ваги, f – порогова функція графу, ι_{ij} – оператор, що враховує ступінь інформованості про стан об'єкту; ζ_{ij} – оператор для врахування ступеню зашумленості даних про стан об'єкту. Процес розрахунку є ітеративним – після завдання початкових станів вершин значення станів перераховуються до тих пір, поки різниця між поточними та попередніми станами не виявиться меншою за деяке задане значення.

З виразу (1) можна зробити висновок, що вираз дозволяє описати процеси в об'єкті моніторингу. Зазначений опис є універсальним та дозволяє описати об'єкт аналізу з урахуванням ієрархічності та індивідуальної специфіки кожного об'єкту моніторингу. При записі виразу (1) в формі багатовимірного часового ряду, процес опису можна привести для динамічної системи. Вираз (1) при побудові математичного опису стану об'єкту моніторингу враховує ступінь інформованості про стан об'єкту та зашумленість даних.

Наступним кроком подальшого дослідження авторів слід вважати визначення деструктивного впливу на об'єкт моніторингу. Сам з цією метою пропонується проведення математичної моделі визначення деструктивного впливу на об'єкт моніторингу.

За визначенням виявлення деструктивного впливу на об'єкт моніторингу включає в себе не тільки ідентифікацію по шаблонам, проте і детектування впливу, що раніше не зустрічався. Разом з тим методи машинного навчання в контексті постановки такого завдання спрямовані всього лише саме на пошук взаємозв'язків і закономірностей функціонування об'єкт моніторингу, знаходженні активності, що схожа на ту, що раніше зустрічалася в навчальній вибірці [354].

Застосування інструментів машинного навчання в готовому вигляді до завдання виявлення деструктивного впливу на об'єкт моніторингу призводить до великої кількості не виявлених впливів та дезорганізації самої об'єкт моніторингу. Насамперед це зумовлено динамічністю ведення бойових дій підрозділами (проведення операцій угрупованнями військ (сил)) та

неоднорідним трафіком, що циркулює в мережі. Крім того, важко відстежити циклічність або сезонність такого обміну [11–15].

Тому для навчання інтелектуальних систем управління процесом ідентифікації об'єктів моніторингу пропонується наступний підхід:

максимально можливий опис множини всіх контрольованих параметрів об'єктів моніторингу;

застосування методів кореляційного аналізу для усунення компонентів, а іноді їх лінійних комбінацій, з близькою до нуля дисперсією;

набір ознак, що залишився, використовується для навчання і перевірки моделі машинного навчання.

Для цього пропонується провести розробку штучної імунної системи, на для виявлення та ідентифікації деструктивного впливу на об'єкти моніторингу.

В даному дослідженні пропонується модель штучної імунної системи на основі еволюційного підходу для ідентифікації деструктивного впливу на об'єкти моніторингу, яка описується:

$$AISEA = \langle D_\tau, D_M, S_A, S_N, G, R, \Psi \rangle, \quad (2)$$

де $D_\tau \subset D$ – набір часових імунних детекторів, D_M – набір імунних детекторів пам'яті, $S_A \subset S$ – навчальна вибірка, що складається з аномальних екземплярів (набір відомих варіантів деструктивного впливу на об'єкти моніторингу), $S_N \subset S$ – тестова множина, що складається з нормальних екземплярів (множина параметрів об'єктів моніторингу), $D = D_\tau \cup D_M$ – набір імунних детекторів, $G = \{G_1, \dots, G_K\}$ – стратегії генетичної оптимізації імунних детекторів, $R: D \times 2^{S_A} \times 2^{S_N} \times G \rightarrow D$ – правило навчання імунних детекторів, S – набір можливих входних об'єктів, $\Psi: D \times S \rightarrow \square_+$ – функція обчислення афінності (правило відповідності) між імунним детектором $d \in D$ та тестовим об'єктом $s \in S$, де $\square_+ = \square \cap [0, +\infty)$.

Кожен імунний детектор $d \in D$ описується як кортеж наступного виду:

$$d = \langle representation, threshold, life_time, state \rangle, \quad (3)$$

де $representation \in \{BitString, RealVector, NeuralNetwork, PetrNet..., \}$ – внутрішнє представлення (внутрішня структура) імунного детектора d , який може бути заданий як двійковий рядок з правилом r -неперервних бітів ($BitString$), дійсний вектор ($RealVector$), штучна нейронна мережа ($NeuralNetwork$), мережа Петрі ($PetriNet$) та ін., $threshold \in \mathbb{R}_+$ – поріг активації імунного детектора d , $life_time \in \mathbb{R}_+$ – термін дії імунного детектора d , $state \in \{immature, semimature, mature, memory\}$ – поточний стан імунного детектора (1, яке може представляти собою незрілий, напівзрілий, зрілий стани або стан, що відповідає детектору пам'яті $\mathbb{R}_+^* = \mathbb{R}_+ \cup \{+\infty\}$).

Загальний підготовчий процес побудови параметрів штучної імунної системи для ідентифікації деструктивного впливу на об'єкти моніторингу може бути описаний таким чином:

1. Визначення корегувального коефіцієнту на ступінь інформованості сили та засоби деструктивного впливу на об'єкти моніторингу. Ступінь інформованості може бути: повна невизначеність, часткова невизначеність, повна обізнаність
2. Вибір внутрішньої структури кожного детектора $d \in D: representation$.
3. Формування навчального набору даних S_A , що містить заздалегідь відібрані „чужі“ об'єкти.
4. Формування тестового набору даних S_N , що містить заздалегідь відібрані „свої“ об'єкти.
5. Вибір стратегії генетичної оптимізації імунних детекторів.
6. Вибір алгоритму навчання R імунних детекторів D в залежності від їх внутрішнього представлення [349].
7. Вибір правила відповідності Ψ між імунним детектором та вхідним об'єктом.

Фактично, запропонована модель штучної імунної системи – це набір імунних детекторів, представлених у вигляді часових детекторів і детекторів пам'яті, в сукупності із заданим алгоритмом їх навчання, який приймає як вхідні

аргументи детектор, що налаштовується d , підмножини двох наборів даних (набору S_A , що складається з „чужих“ об’єктів, та набору S_N , що складається зі „своїх“ об’єктів), а також стратегію генетичної оптимізації.

Причому набір даних призначається для першого попереднього налаштування імунних детекторів, а роль набору даних S_N полягає в фільтрації навчених детекторів. Стратегії генетичної оптимізації імунних детекторів включають деякий набір генетичних операторів (кросоверу, мутації, інверсії) та їх комбінацій для зміни параметрів імунного детектора після його клонування. Правило навчання імунних детекторів є двокроковою процедурою. На першому кроці імунні детектори навчаються виключно на елементах набору даних S_A та піддаються клональній селекції, в процесі якої створені копії імунних детекторів мутують згідно з обраною стратегією. $G' \in G$. З набору детекторів, представленого виробленим нащадків і вихідним детектором, відбираються ті детектори, які є найбільш придатними по відношенню до елементів набору згідно з обраним правилом відповідності Ψ .

Ця фаза повторюється кілька разів для формування напівзрілих детекторів. На другій фазі навчання ці детектори перевіряються на відповідність „своїм“ об’єктам: ті з них, які помилково активуються, знищуються, заново ініціалізуються та навчаються. В рамках цієї моделі кожен імунний детектор піддається кільком етапам диференціювання. На початку свого розвитку кожен детектор ініціалізується довільним чином відповідно до свого внутрішнього уявлення. У процесі функціонування штучної імунної системи часові детектори здійснюють запис параметрів деструктивного впливу на об’єкти моніторингу в базу даних, що оновлюється S_{A^*} у разі виявлення.

Усі імунні детектори, крім детекторів пам’яті, мають кінцевий термін життя. Якщо протягом даного терміну життя вони не виявили жодного деструктивного впливу на об’єкти моніторингу, вони піддаються повторному навчанню на розширеному наборі даних, що містить елементи S_A та S_{A^*} .

У той же час, якщо імунний детектор розпізнав деструктивний вплив на об’єкти моніторингу, його термін життя збільшується. На відміну від них,

детектори пам'яті мають нескінченний термін життя і не беруть участі в наповненні оновлюваного набору деструктивного впливу.

Для виявлення кожного класу деструктивного впливу $c \in C$ виділяється кілька імунних детекторів, що об'єднуються в клас детекторів $D_{\varsigma(C)}$, $\left(\bigcup_{c \in C} D_{\varsigma(C)} = D\right)$. Кожен із детекторів $d \in D_{\varsigma(C)}$ використовує глибоке навчання, запропоноване в роботі [349] на різних випадкових підвиборках множин S_A та $S_N - S_A^{(d)}$, і $S_A^{(d)}$ які, можливо, містять деякі дублюючі та переупорядковані об'єкти з вихідних наборів. Тим самим досягається різноманітність імунних детекторів усередині $D_{\varsigma(C)}$.

Відповідно під q -ою групою детекторів розуміється набір детекторів $D_{\varsigma(C)}$ $\left(\bigcup_{q=1}^m D^{(q)} = D\right)$, m – число класів деструктивного впливу на об'єкти моніторингу), які повністю покривають задану множину класів атак. Якщо детектор d реагує на будь-який елемент $s' \in S_N^{(d)}$ тобто якщо $\exists s' \in S_N^{(d)} \Psi(d, s') > \min_{s \in S_A^{(d)}} \Psi(d, s)$, то детектор d піддається апоптозу та заміні новим довільно згенерованим детектором. Для кожного класу деструктивного впливу $c \in C$ визначається рівно один детектор пам'яті $d_m^{(c)}$ – детектор, який задовольняє вимозі максимальної пристосованості до розпізнавання об'єктів $S_A^{(d_m^{(c)})} \cap C$. Таким чином, набір імунних детекторів пам'яті може бути визначений таким чином:

$$D_M = \bigcup_{c \in C} \left\{ \arg \max_{d \in D_{\varsigma(C)}} \left(\frac{\sum_{s \in S_A^{(d)} \cap C} \Psi(d, s)}{\#(S_A^{(d)} \cap C)} \right) \right\}. \quad (4)$$

Набір часових імунних детекторів визначається наступним чином:

$$D_\tau = D / D_M. \quad (5)$$

Поріг активації імунного детектора $d \in D$, навченого на наборах $S_A^{(d)}$ та $S_N^{(d)}$ обчислюється наступним чином:

$$threshold = \frac{\overbrace{\min_{s \in S_A^{(d)}} \Psi(d, s)}^{h_d^-} + \overbrace{\min_{s \in S_N^{(d)}} \Psi(d, s)}^{h_d^+}}{2}. \quad (6)$$

Зазначена формула застосовується для обчислення порогу активації тільки тих детекторів d , які після навчання на множині аномальних даних $S_A^{(d)}$ не мають помилкових спрацьовувань на множині нормальних даних $S_A^{(d)}$, тобто $\forall s' \in S_N^{(d)} \Psi(d, s') < h_d^-$. Якщо така умова виконується, то з'являється додатковий проміжок, рівний величині $h_d = h_d^- - h_d^+ > 0$, і тим самим виникає можливість „зрушити“ граничне значення h_d^- , що забезпечує реагування детектора d на будь-який „чужий“ об'єкт $s \in S_A^{(d)}$ на величину $\frac{h_d^- - h_d^+}{2}$ в сторону h_d^+

$$threshold = h_d^- - \frac{h_d^- - h_d^+}{2} = \frac{h_d^- + h_d^+}{2}.$$

Виявлення деструктивного впливу на об'єкти моніторингу $s \in S$ за допомогою розглянутої моделі здійснюється наступним чином:

1. Визначення корегувального коефіцієнту на ступінь інформованості сили та засоби деструктивного впливу на об'єкти моніторингу.

2. Обчислення для кожного імунного детектора $d \in D$ значення його активації $a_d = \Psi(d, s) - threshold$. Вважається, що якщо $a_d \geq 0$, то детектор d є активованим, інакше відповідний детектор не реагує на вхідний об'єкт.

3. Мажоритарне голосування всередині кожного класу детекторів $D_{\zeta(C)}$

$$\underbrace{\sum_{d \in D_{\zeta(C)}} [a_d \geq 0]}_{A_c} > \underbrace{\sum_{d \in D_{\zeta(C)}} [a_d < 0]}_{B_c = \#D_{\zeta(C)} - A_c}, \text{ то } s \text{ розпізнається як „чужий“ об'єкт. Якщо } A_c < B_c, \text{ то}$$

s розпізнається як „свій“ об'єкт. В разі наявності конфліктів, тобто $A_c = B_c$, s класифікується як „чужий“ об'єкт, якщо $a_{d_m^{(C)}} \geq 0$, та s класифікується як „свій“ об'єкт, якщо $a_{d_m^{(C)}} < 0$, де $d_m^{(C)} \in D_M \cap D_{\zeta(C)}$, $d_m^{(C)}$ – детектор пам'яті, навчений для розпізнавання „свого“ об'єкту та „чужого“ об'єкта з класу C .

4. Формування множини класів імунних детекторів, що активувалися $\{D_{\varepsilon(C')}\}_{C' \in c}$ які розпізнають вхідний об'єкт s як „чужий“ об'єкт, де $C^* = \left\{ C' \mid C' \in c \wedge \left(A_{C'} > B_{C'} \vee \left(A_{C'} = B_{C'} \wedge a_{d_m^{(C')}} \geq 0 \right) \right) \right\} \subset c$.

5. Визначення класу об'єкта s . Якщо $C^* = \emptyset$, то об'єкт s належить до класу „своїх“ об'єктів. Якщо $E_{c^*} = \max_{C' \in C^*} A_{C'}$ досягається в одній єдиній точці, то клас об'єкта $\arg E_{c^*}$, інакше клас об'єкта s - це $\arg \max_{C' \in \{\arg E_{c^*}\}} \sum_{d \in D_{\varepsilon}(C')} \times [a_d \geq 0]$.

Даний алгоритм заснований на порівнянні величини афінності імунних детекторів з їх індивідуально налаштованими порогоми активації та врахуванні однакових голосів, отриманих від більшої частини детекторів. У разі виникнення конфліктів при розрізненні між нормальним та аномальним класом (крок 3) вирішальний голос віддається детектору пам'яті. Якщо ж після цього зберігається конфлікт лише на рівні групи детекторів, то береться до уваги сума величин афінності, що саме активувалися у відповідь на даний стимул (вхідний об'єкт) імунних детекторів (крок 5).

Вхідний об'єкт є „своїм“ тоді й лише тоді, коли $\forall C \in c \ A_C < B_C \vee \left(A_C = B_C \wedge a_{d_m^{(C)}} < 0 \right)$.

Як основу для даної моделі використовувалися модель з життєвим циклом (M_1), запропонована в [355], і модель з бібліотекою генів (M_2) [356], запропонована в [357], модель AISEA (M_3) [358] та розроблена модель (M_4), що була доповнена низкою удосконалень, а саме:

врахуванням типу невизначеності про об'єкти моніторингу;

удосконаленням алгоритмом навчання імунних детекторів;

механізмом автоматичного підбору їх порогу активації, а також процедурою вирішення конфліктних випадків класифікації одного об'єкта за допомогою імунних детекторів пам'яті.

Порівняння цих чотирьох моделей наведено в таблиці 1. Знаком “+” відмічені характеристики, які притаманні відповідній моделі, знак “-” означає відсутність підтримки цієї особливості моделі.

Таблиця 1

Порівняння імунних моделей для виявлення деструктивного впливу на
об'єкти моніторингу

Імунна модель	Незалежність від внутрішньої структури імунного детектору	Наявність клональної селекції та генетичної оптимізації	Наявність негативного відбору	Врахування типу невизначеності РЕО	Автоматичний підбір порога активації імунного детектора	Наявність механізмів навчання	Наявність детекторів пам'яті	Наявність життєвого циклу лімфоцитів	Динамічне перенавчання	Навчання детекторів на нових даних в процесі функціонування	Розподіл імунних детекторів	Підтримка мультикласового виявлення деструктивного впливу на СРЗ	Автономність імунної системи (робота без втручання оператора)	Наявність сукупності процедур обробки різних даних
M_1	–	–	+	–	–	–	+	+	+	+	–	–	–	–
M_2	–	+	+	–	–	–	+	–	+	+	+	–	–	–
M_3	+	+	+	–	+	+	+	+	+	+	–	–	+	–
M_4	+	+	+	+	+	+	+	+	+	+	+	+	+	+

Результати порівняльного аналізу, що наведені в таблиці 1 дозволяють зробити висновок про перевагу зазначеної моделі у порівнянні з відомими.

Зазначений підхід запропоновано використовувати в ході урегулювання воєнних конфліктів. Це дозволить підвищити оперативність обробки та передачі даних.

Розроблена модель M_4 є універсальною по відношенню до внутрішнього представлення імунних детекторів, у той час як модель M_1 орієнтована на використання бітових рядків як імунні детектори, а модель M_2 – на кластерну дискретизацію полів (генів) імунних детекторів. У моделі M_1 відсутня можливість клональної селекції детекторів, а також не передбачені оператори генетичної мутації. У всіх із представлених моделей у контексті виявлення деструктивного впливу на об'єкти моніторингу використовується динамічно оновлювана популяція імунних детекторів, що дозволяє штучній імунній системі адаптуватися до радіоелектронної обстановки, що змінюються, в режимі її

функціонування. Однак тут для моделі M_1 вводиться додаткове обмеження: після активації детектора в результаті накопичення достатньої кількості збігів з антигенами повинен бути згенерований зовнішній по відношенню до системи сигнал зі сторони адміністратора (сигнал костимуляції). Такий сигнал дозволить активованому детектору отримати шанс для вступу в пул детекторів пам'яті та можливого подальшого безперервного аналізу радіоелектронної обстановки. Для моделі M_2 характерний розподілений механізм генерації детекторів: детектори, вигідні з точки зору розпізнавання аномалій, можуть бути розмножені інші мережеві вузли підвищення загальної продуктивності системи; ця властивість відсутня у моделях M_1 та M_3 .

Проте в розробленій моделі M_4 додатково:

- враховуються тип невизначеності про наявні можливості системи деструктивного впливу на об'єкти моніторингу;
- використовується удосконалена сукупність процедур обробки різнотипних даних;
- використовується удосконалена процедура навчання;
- використовується механізм розв'язання конфліктних випадків класифікації;
- процедурою автоматичного обчислення порога активації імунних детекторів, а також універсальністю структури їхнього представлення;
- постійним оновлення імунних детекторів протягом різних етапів дозрівання (життєвого циклу) та їх перенавчання з використанням набору деструктивного впливу на об'єкти моніторингу, що розширюється.

Обмеженнями зазначеного дослідження слід вважати:

- врахування часових обмежень на передачу конкретного типу повідомлення (формалізованого донесення);
- наявність первинної бази про об'єкт моніторингу;
- обмеження щодо якості каналів передачі даних.

Алгоритм реалізації методу оцінки в інтелектуальних системах підтримки прийняття рішень

Метод оцінки в інтелектуальних системах підтримки прийняття рішень складається з наступної послідовності дій (рис. 1).

Алгоритм реалізації запропонованого методу складається з наступної послідовності дій.

1. *Введення вихідних даних.* На даному етапі відбувається введення вихідних даних про стан об'єкту моніторингу. Визначається кількість джерел технічних засобів моніторингу, тип вихідних даних та їх обсяг.

2. *Визначення ступеня невизначеності вихідних даних.* На даному етапі визначається ступінь невизначеності вихідних даних на підставі попередніх досліджень авторів. Ступінь невизначеності вихідних даних наступна: повна невизначеність; часткова невизначеність та повна обізнаність [345, 366].

3. *Побудова дерева класифікаторів.*

Зазначений етап методу може бути охарактеризований як підготовчий, він містить у собі вибір структури окремих бінарних класифікаторів (детекторів):

розмірності та числа шарів, параметрів і алгоритмів навчання, типів функцій активації, функцій належності та ядерних функцій [24–28].

Для кожного детектора складається набір навчальних правил. Задаючи різну сукупність таких наборів правил, можна сформувати групу детекторів, кожний з яких побудовано на основі штучної нейронної мережі, що еволюціонує. Детектори усередині кожної такої групи поєднуються в класифікатор на основі підходів один-до-усіх (one-vs-all), один-до-одного (one-vs-one) або їх різних похідних варіацій [28–35].

У першому підході кожний детектор $F_{jk}^{(k)}: \square \rightarrow \{0,1\} (k=1,...,m)$ навчається на даних $\{x_l, [c_l = k]\}_{l=1}^M$, і функціонування групи детекторів $F_{jk}^{(k)}$ описується за допомогою принципу, що виключає:

$$F_j^{(i)}(z) = \begin{cases} \{0\}, \text{ якщо } \forall k \in \{1, ..., m\} F_{jk}^{(i)}(z) = 0 \\ \{k \mid F_{jk}^{(i)}(z) = 1\}_{k=1}^m, \text{ інакше} \end{cases}, \quad (7)$$

У другому підході кожний з $\tilde{N}_{m+1}^2 = \frac{(m+1) \cdot m}{2}$ детекторів $F_{jk_0k_1}^{(k)}$ навчається на множині об'єктів, що належать тільки двом класам з мітками k_0 і k_1 , $-\{(x_l, 0 | \bar{c}_l = k_0)\}_{l=1}^M \cup \{(x_l, 1 | \bar{c}_l = k_1)\}_{l=1}^M, 0 \leq k_0 < k_1 \leq m$ та функціонування групи детекторів $F_j^{(i)}$ задається за допомогою голосування max-wins:

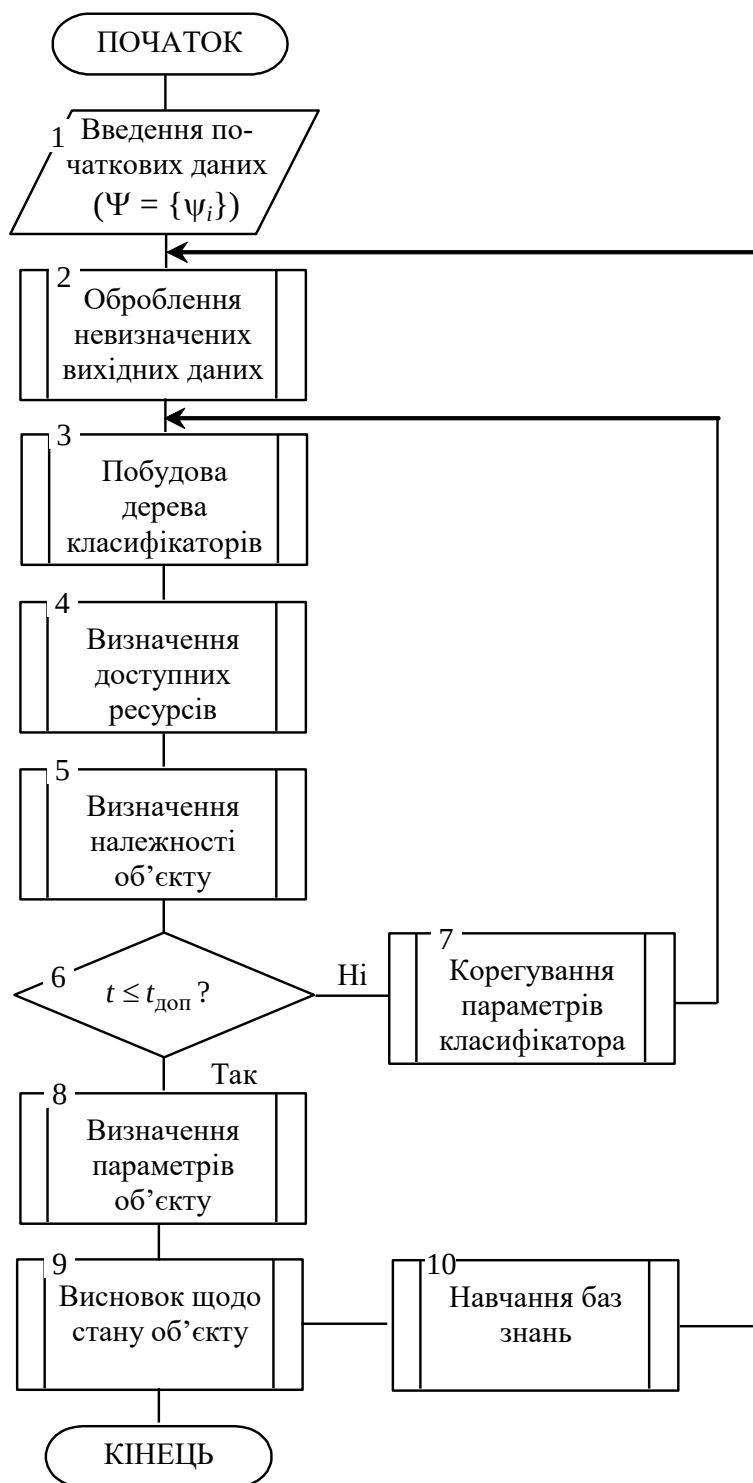


Рис. 1. Алгоритм реалізації методу аналізу стану об'єкту

$$F_j^{(i)} = \left\{ \arg \max_{\bar{c} \in \{0, \dots, m\}} \sum_{k=\bar{c}+1}^m [F_{j\bar{c}k}^{(i)}(z) = 0] + \sum_{k=0}^{\bar{c}-1} [F_{j\bar{c}k}^{(i)}(z) = 1] \right\}. \quad (8)$$

4. Визначення доступних апаратних обчислювальних ресурсів.

На даному етапі визначаються доступні апаратні обчислювальні ресурси мережі. На підставі чого визначаються можливі варіанти класифікації: бінарне класифікаційне дерево, генетичний алгоритм, нечіткі когнітивні моделі та ациклічний граф.

У таблиці 2 наведені характеристики розглянутих схем об'єднання детекторів у багатокласову модель, призначену для співвіднесення вхідного об'єкта однієї або декільком з $(m+1)$ міток класів.

Таблиця 2

Характеристики схем об'єднання детекторів

Схема об'єднання	Число детекторів, що підлягає навчанню	Мінімальне число детекторів, задіяних при класифікації об'єктів	Максимальне число детекторів, задіяних при класифікації об'єктів
один-до-усіх	m	m	m
один-до-одного	$\frac{(m+1) \cdot m}{2}$	$\frac{(m+1) \cdot m}{2}$	$\frac{(m+1) \cdot m}{2}$
Класифікаційне бінарне дерево	m	1	m
Спрямований ациклічний граф	$\frac{(m+1) \cdot m}{2}$	m	m
Нечітка когнітивна модель	$(m \cdot x) \cdot m$	$(m \cdot x)$	$(m \cdot x) \cdot m$
Генетичний алгоритм	$(m \cdot x)$	m	$(m \cdot x)$

5. Визначення належності об'єкту моніторингу до певного класу

У якості однієї з похідних варіацій попередніх підходів для комбінування детекторів може бути згадане класифікаційне бінарне дерево [368]. Формально така структура задається рекурсивно в такий спосіб:

$$CBT_{\mu} = \begin{cases} \langle F_{jL_{\mu}R_{\mu}}^{(i)}, CBT_{L_{\mu}}, CBT_{R_{\mu}} \rangle, & \text{якщо } \# \mu \geq 2 \\ \mu, & \text{якщо } \# \mu = 1. \end{cases} \quad (9)$$

де $\mu = \{0, \dots, m\}$ – вихідний набір міток класів, L_{μ} Ш μ – довільно згенерована або визначене підмножина; $\mu(\#L_{\mu} < \# \mu), R_{\mu} = \mu \setminus L_{\mu}$ – ліве класифікаційне піддерево, $CBT_{R_{\mu}}$ – праве класифікаційне піддерево, $F_{jL_{\mu}R_{\mu}}^{(i)}$ – вузловий детектор, навчений на елементах множини $\{(x_l, 0) | \bar{c}_l \in L_{\mu}\}_{l=1}^M \cup \{(x_l, 1) | \bar{c}_l \in R_{\mu}\}_{l=1}^M$.

Вихідний результат детектора настроюється таким чином, щоб він дорівнював 0, якщо вхідний об'єкт x_l належить класу з міткою $\bar{c}_l \in L_{\mu}$, і 1, якщо об'єкт x_l належить класу з міткою $\bar{c}_l \in L_{\mu}$

Тому функціонування групи детекторів $F_j^{(i)}$, представлених у вигляді вузлів такого дерева, описується за допомогою рекурсивної функції $\phi_j^{(i)}$, що задає послідовну дихотомію множини μ :

$$F_j^{(i)} = \phi_j^{(i)}(\mu, z),$$

$$\phi_j^{(i)}(\mu, z) = \begin{cases} \mu, & \text{якщо } \# \mu = 1 \\ \phi_j^{(i)}(L_{\mu}, z) & \text{якщо } \# \mu \geq 2 \wedge F_{jL_{\mu}R_{\mu}}^{(i)}(z) = 0 \\ \phi_j^{(i)}(R_{\mu}, z) & \text{якщо } \# \mu \geq 2 \wedge F_{jL_{\mu}R_{\mu}}^{(i)}(z) = 1. \end{cases} \quad (10)$$

Застосування функції $\phi_j^{(i)}$ до вихідного набору міток класів і об'єкту моніторингу дозволяє здійснювати однозначний пошук мітки класу цього об'єкту. Це пояснюється тим, що оскільки під час спуску вниз по класифікаційному дереву відбувається диз'юнктивне розбиття множини міток класів. Після досягнення та спрацьовування термінального детектора залишається тільки одна можлива мітка для класифікації вхідного об'єкта z у якості вихідного результату $F_j^{(i)}$. Тому для класифікаційного дерева неможливі

конфліктні випадки при класифікації об'єктів, які можуть мати місце для двох інших підходів комбінування.

Іншим підходом є спрямований ациклічний граф, який організує $C_{m+1}^2 = \frac{(m+1) \cdot m}{2}$ детекторів у зв'язну динамічну структуру, яка може бути задана наступною формулою:

$$DAG_{\mu} = \begin{cases} \left\langle F_{j\mu k_0 k_1}^{(i)}, DAG_{\mu \setminus \{k_0\}}, DAG_{\mu \setminus \{k_1\}} \right\rangle, & \text{якщо } \# \mu \geq 2, \text{ де } k_0 \in \mu, k_1 \in \mu, \\ \mu, & \text{якщо } \# \mu = 1. \end{cases} \quad (11)$$

Тут, як і в підході один-до-одного, кожний вузловий детектор $F_{j\mu k_0 k_1}^{(i)}$ навчається на елементах $\left\{ (x_l, 0 | \bar{c}_l = k_0) \right\}_{l=1}^M \cup \left\{ (x_l, 1 | \bar{c}_l = k_1) \right\}_{l=1}^M (k_0 < k_1)$. Обхід розглянутого графа виконується за допомогою рекурсивної функції $\xi_j^{(i)}$, що задає заелементне „відщиплення“ від множини μ :

$$F_j^{(i)} = \xi_j^{(i)}(\mu, z),$$

$$\xi_j^{(i)}(\mu, z) = \begin{cases} \mu, & \text{якщо } \# \mu = 1 \\ \xi_j^{(i)}(\mu \setminus \{k_1\}, z), & \text{якщо } \# \mu \geq 2 \wedge F_{j\mu k_0 k_1}^{(i)}(z) = 0 \\ \xi_j^{(i)}(\mu \setminus \{k_0\}, z), & \text{якщо } \# \mu \geq 2 \wedge F_{j\mu k_0 k_1}^{(i)}(z) = 1. \end{cases} \quad (12)$$

Якщо детектор $F_{j\mu k_0 k_1}^{(i)}$ голосує за k_0 -ий клас для об'єкта z , тобто $F_{j\mu k_0 k_1}^{(i)}(z) = 0$, то з множини μ видаляється мітка k_1 як свідомо невірна, а якщо ні, то виключається мітка k_0 . Процес повторюється доти, поки множина μ не вироджується в одноелементне.

З розглянутих шести схем тільки одна, а саме класифікаційне бінарне дерево, має змінне число детекторів, які можуть використовуватися в процесі класифікації об'єктів.

Мінімальне значення досягається, коли активується детектор $F_{jL_{\mu} R_{\mu}}^{(i)}$, розташований у корені дерева та навчений для розпізнавання тільки одного класу об'єктів серед усіх інших, і $F_{jL_{\mu} R_{\mu}}^{(i)}(z) = 0 (F_{jL_{\mu} R_{\mu}}^{(i)}(z) = 1)$ тобто коли $\# L_{\mu} = 1 (\# R_{\mu} = 1)$. Максимальне значення досягається, коли дерево представляється послідовним списком і активується найбільш вилучений у ньому детектор.

У випадку збалансованого дерева цей показник може становити величину $\lfloor \log_2(m+1) \rfloor$ або $\lceil \log_2(m+1) \rceil$. Кожний класифікатор $F^{(i)} (i=1, \dots, P)$ містить q_i груп $F_j^{(i)} (j=1, \dots, q_i)$, кожна з яких поєднує m детекторів $F_{jk}^{(i)} (k=1, \dots, m)$ за допомогою підходу один-до-усіх. Кожна із груп детекторів $F_j^{(i)}$ навчається на різних випадкових вибірках, які можуть включати повторювані та переупорядковані елементи з вихідного навчального набору $\Upsilon_{\chi_c^{(LS)}}$. Об'єднання груп $F_j^{(i)}$ у класифікатор $F^{(i)}$ здійснюється на основі гібридного правила, що представляє собою суміш голосування більшістю й голосування max-wins:

$$F^{(i)}(z) = \left\{ \bar{c} \left| \underbrace{\sum_{j=1}^{q_i} [\bar{c} \in F_j^{(i)}(z)]}_{\Xi_i(\bar{c})} > \frac{1}{2} \cdot q_i \wedge \Xi_i(\bar{c}) = \max_{\bar{c}' \in \{0, \dots, m\}} \Xi_i(\bar{c}') \right. \right\}_{\bar{c}'=0}^m. \quad (13)$$

В даній формулі за рахунок вимоги $\Xi_i(\bar{c}) > \frac{1}{2} \cdot q_i$ класифікатор $F^{(i)}$ стає нездатним вирішувати конфлікти, які виникають за умови $\# \left\{ \bar{c} \left| \Xi_i(\bar{c}) = \frac{1}{2} \cdot q_i \wedge \Xi_i(\bar{c}) = \max_{\bar{c}' \in \{0, \dots, m\}} \Xi_i(\bar{c}') \right. \right\}_{\bar{c}'=0}^m = 2$ (у цьому випадку виходом класифікатора є порожня множина $\emptyset$).

У процесі роботи інтерпретатора перевіряється коректність оброблюваних даних і ініціалізуються поля об'єктів усередині дерева класифікаторів. За рахунок використання такої структури в рамках запропонованого методу стає можливим будувати багаторівневі схеми.

Даний метод має розподілену архітектуру, у яких збір даних здійснюється вторинними вузлами-сенсорами (технічними засобами розвідки), а вся обробка агрегованих потоків даних виконується на централізованому сервері.

6. Визначення параметрів об'єкту відповідного класу

Зазначений етап методу, виконуваний на стороні сенсорів (технічних засобів розвідки), полягає в складанні необроблених розвідувальних відомостей у класифікаційні блоки, виділенні їх параметрів і виконанні аналізу з використанням декількох паралельних алгоритмів шаблонного пошуку.

Суть процедури полягає в розбивці заданого часового інтервалу $\Delta_0^{(L)} = [0, L]$ довжиною L , протягом якого проводиться безперервне спостереження за рядом параметрів, на трохи більш дрібних інтервалах $\Delta_0^{(L')}, \Delta_{\delta}^{(L')}, \dots, \Delta_{\delta(k-1)}^{(L')}$ однакової довжини $0 < L' \leq L$, початок кожного з яких має зсув $0 < \delta \leq L'$ відносно початку попереднього інтервалу. Причому $\bigcup_{i=0}^{k-1} \Delta_{\delta-i}^{L'} \subseteq \Delta_0^{(L)}$ й $\bigcup_{i=0}^k \Delta_{\delta-i}^{L'} \supseteq \Delta_0^{(L)}$ тому $k = 1 + \left\lceil \frac{L-L'}{\delta} \right\rceil$. Протягом проміжків часу $\Delta_0^{(L')}, \dots, \Delta_{\delta(k-1)}^{(L')}$ відбувається фіксація значень $\omega_0, \dots, \omega_{k-1}$ параметрів, і їх середня величина (інтенсивність) та у рамках часового вікна довжини L' розраховується по формулі $\bar{\omega} = \frac{1}{k} \cdot \sum_{i=0}^{k-1} \omega_i$.

У дослідженні використовувався інтервал зі значенням параметра L , рівним п'яти секундам. Довжина інтервалу, що L' згладжує, була обрана рівній однієї секунді. Зсув δ був установлений у півсекунди. подібний підхід дозволяє усунути рідкі по частоті й випадкові мережні сплески й тим самим знизити число неправильних спрацьовувань.

7. Попередня обробка даних про об'єкт аналізу.

Перед безпосереднім навчанням детекторів виконується попередня обробка даних параметрів для зменшення ефекту їх сильної мінливості.

Багато методів, включаючи нейронні мережі та метод головних компонентів, чутливі до такого роду флуктуаціям і вимагають, щоб усі ознаки оброблюваних векторів мали однаковий масштаб.

7.1. Нормалізація компонентів вектору.

Перший крок попередньої обробки кожного компонента x_{ij} вектору

$x_i \in \{x_k\}_{k=1}^M$ включає його нормалізацію за допомогою функції $f(x_{ij}) = \frac{x_{ij} - x_j^{(\min)}}{x_j^{(\max)} - x_j^{(\min)}}$ (у

випадку $x_j^{(\max)} = x_j^{(\min)}$ можна вважати $f(x_{ij}) = 0$), де $x_j^{(\min)} = \min_{1 \leq i \leq M} x_{ij}$ та $x_j^{(\max)} = \max_{1 \leq i \leq M} x_{ij}$

7.2 Мінімізація простору ознак.

Зменшення числа значимих ознак, яке досягається за допомогою методу головних компонентів [26–30], описуваного як послідовність наступних кроків:

7.2.1 Обчислення математичного очікування випадкового вектору, представленого в цьому випадку у вигляді елементів набору навчальних даних:

$$\left\{ x_i = \left\{ x_{ij} \right\}_{j=1}^n \right\}_{i=1}^M : \\ \bar{x} = \left(\frac{1}{M} \cdot \sum_{i=1}^M x_{i1}, \dots, \frac{1}{M} \cdot \sum_{i=1}^M x_{in} \right)^T. \quad (14)$$

7.2.2. Формування елементів незміщеної теоретичної коваріаційної матриці:

$$\Sigma = (\sigma_{ij})_{\substack{i=1, \dots, n \\ j=1, \dots, n}} : \\ \sigma_{ij} = \frac{1}{M-1} \cdot \sum_{k=1}^M (x_{ki} - \bar{x}_i) \cdot (x_{kj} - \bar{x}_j). \quad (15)$$

7.2.3. Знаходження власних чисел $\{\lambda_i\}_{i=1}^n$ і власних векторів $\{v_i\}_{i=1}^n$ матриці Σ як корінь рівнянь (із цією метою використовувався метод обертань Якобі):

$$\begin{cases} \det(\Sigma - \lambda \cdot \mathbf{I}) = 0 \\ (\Sigma - \lambda \cdot \mathbf{I}) \cdot v = 0, \end{cases} \quad (16)$$

де $\mathbf{I}$ — одинична матриця розміром $n \times n$.

7.2.4. Ранжування власних чисел $\{\lambda_i\}_{i=1}^n$ у порядку їх убутання й відповідних їм власних векторів $\{v_i\}_{i=1}^n$:

$$\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n \geq 0. \quad (17)$$

7.2.5. Відбір необхідного числа $\hat{n} \leq n$ головних компонентів:

$$\hat{n} = \min \left\{ j \mid \varsigma(j) \geq \varepsilon \right\}_{j=1}^n, \quad (18)$$

де $\varsigma(j) = \frac{\sum_{i=1}^j \lambda_i}{\sum_{i=1}^n \lambda_i}$ - захід інформативності [344], $0 \leq \varepsilon \leq 1$ – експертно обирає

величина.

7.2.6 Центрування вектору ознак $z: z_c = z - \bar{x}$.

7.2.7 Проектування центрованого вектору ознак z_c у нову систему координат, що задається ортонормованими векторами $\{v_i\}_{i=1}^{\hat{n}}$

$$y = (y_1, \dots, y_{\hat{n}})^T = (v_1, \dots, v_{\hat{n}})^T \cdot z_c,$$

$y_i = v_i^T \cdot z_c$ називається i -ою головною компонентною вектору z .

8. Ієрархічний обхід дерева класифікаторів по ширині.

Зазначений етап методу з погляду обчислювальних ресурсів є найбільш трудомістким і складається з наступних рекурсивно повторюваних послідовностей дій: обчислення залежностей поточного класифікатора, формування вхідних сигналів для поточного класифікатора та навчання поточного класифікатора.

Навчання кожного класифікатора породжує запит на навчання класифікаторів, що внизу, зазначених у списку його залежностей, і генерацію їх вихідних даних для формування вхідних даних класифікатора, що вище.

Наслідком використовуваного в такий спосіб каскадного навчання є можливість “ледачого” завантаження класифікаторів: у навчанні й розпізнаванні беруть участь тільки ті класифікатори, які зустрічаються в списку залежностей класифікатора, відповідального за формування загального рішення в колективі класифікаційних правил.

Ця властивість є особливо вигідною при розборі динамічних правил навчання класифікаторів, тобто таких правил, від успішного або неуспішного спрацьовування яких залежить ініціалізація іншого правила. Зокрема, це характерно для класифікаційного дерева, коли правила є вкладеними один у одного.

Метод надає можливість будувати багаторівневі схеми з довільною вкладеністю класифікаторів один одного і їх “лінивим” підключенням у процесі аналізу вхідного вектору.

Проведено моделювання роботи методу пошуку рішень відповідно до алгоритму на рис. 1 та виразів (1)–(18). Проведено моделювання роботи запропонованої методу оцінки в програмному середовищі MathCad 14 (США). В якості задачі, що вирішувалася при проведенні моделювання була оцінка елементів оперативної обстановки угруповання військ (сил).

Вихідні дані для оцінки стану оперативної обстановки з використанням запропонованого методу:

– кількість джерел інформації, про стан об'єкту моніторингу – 3 (засоби радіомоніторингу, засоби дистанційного зондування землі та безпілотні літальні апарати) Для спрощення моделювання було взято однакову кількість кожного засобу – по 4 засоби;

– кількість інформаційних ознак по яким відбувається визначення стану об'єкту моніторингу – 12. До таких параметрів відносяться: належність, тип організаційно-штатного формування, пріоритетність, мінімальна ширина по фронту, максимальна ширина по фронту. Також враховується кількість особового складу, мінімальна глибина по флангу, максимальна глибина по флангу, загальна чисельність особового складу, кількість зразків ОБТ, кількість типів зразків ОБТ та кількість засобів зв'язку);

– варіанти організаційно-штатних формувань – рота, батальйон, бригада.

Позначимо, які параметри для кожного типу операторів розглядалися. Метод був апробована при пропорційній селекції (обсяг 18 %); рекомбінації: середня. Щоб визначити найбільш ефективну комбінацію налаштувань для кожної окремої розглянутої схеми необхідно всі інші параметри пошуку залишити однаковими. Обсяг популяції був обраний рівним 50, число популяцій – 50. Зазначені дані взяті відповідно до орієнтовної чисельності командних пунктів оперативно-тактичного угруповання військ (сил). Порівняння алгоритмів здійснюється за критерієм придатності отриманих рішень. Число незалежних запусків у експериментах – 100. Швидкість оцінювалася як середнє покоління, на якому алгоритм знаходить глобальний оптимум.

Порівнювалися кілька різних оптимізаційних алгоритмів вирішення поставленого екстремальної задачі (16). Серед них: класичний бінарний генетичний алгоритм; дійсний генетичний алгоритм; запропонований метод та генетичний алгоритм з алгоритмом налаштування Population-Level Dynamic Probabilities (PDP). При цьому кількість обчислень цільової функції для роботи генетичних алгоритмів було вибрано рівним числу вимірювань цільової функції, в циклах яких використовувалося локального поліпшення [367].

В табл. 3 наведені результати порівняння для запропонованого методу та

відомих при пошуку в одному напрямку, пошуку в двох напрямках та пошуку в трьох напрямках.

Таблиця 3

Часові та системні витрати навчання та тестування

Індикатори			Схеми комбінування			
			one-vs-all	one-vs-one	CBT	DAG
Навчання: кількість навчальних екземплярів — 8000						
Однопоточковий режим	Час (сек.)	min	11235.000	4569.000	4679.000	4660.000
		max	12 228.000	5218.000	5092.000	6148.000
		avg	11822.000	4967.333	4880.000	5217.333
	Завантаженість центрального процесора (%)	min	92.700	92.700	92.700	92.700
		max	100.000	100.000	100.000	100.000
		avg	99.215	99.194	99.219	99.220
	Завантаження оперативної пам'яті (%)	min	62.496	64.309	62.492	64.324
		max	92.352	98.996	94.219	99.563
		avg	95.487	97.688	99.067	97.435
багатонаправлений режим	Час (сек.)	min	1782.000	754.000	900.000	743.000
		max	6996.000	875.000	1228.000	880.000
		avg	2767.429	823.429	1059.143	831.571
	Завантаженість центрального процесора (%)	min	92.700	92.700	92.700	92.800
		max	800.000	800.000	800.000	800.000
		avg	455.184	579.644	445.293	610.500
	Завантаження оперативної пам'яті (%)	min	62.496	64.305	62.504	64.324
		max	83.867	66.070	68.820	66.090
		avg	263.452	480.843	248.808	480.758

Тестування: кількість тестовий екземплярів — 16000						
Однопотоковий режим	Час (сек.)	min	383.000	985.000	176.000	296.000
		max	387.000	1019.000	182.000	301.000
		avg	384.667	996.667	179.000	298.000
	Завантаженість центрального процесора (%)	min	92.800	92.800	92.700	92.800
		max	100.000	100.000	100.000	100.000
		avg	99.197	99.177	99.123	99.212
	Завантаження оперативної пам'яті (%)	min	63.656	65.289	63.395	65.348
		max	88.957	90.816	88.125	90.500
		avg	82.348	84.368	79.987	83.629
багатонаправлений режим	Час (сек.)	min	644.000	1764.000	288.000	563.000
		max	701.000	1789.000	324.000	576.000
		avg	660.857	1776.857	308.286	569.429
	Завантаженість центрального процесора (%)	min	86.400	86.300	59.800	59.700
		max	93.000	96.100	93.000	72.800
		avg	113.028	119,584	99.806	100.081
	Завантаження оперативної пам'яті (%)	min	64.477	66.301	63.934	67.000
		max	91.031	92.215	88.473	91.250
		avg	82.987	84.877	80.397	84.525

За результати аналізу даних, що наведені в табл. 3 видно, що запропонований метод має прийнятну обчислювальну складність. Запропонований метод дозволяє отримувати адекватні рішення при складній ієрархічній структурі об'єкту моніторингу. Ефективність запропонованого методу в середньому складає від 12 до 20 % при різних схемах комбінування.

Основними перевагами запропонованого методу оцінки є:

- має гнучку ієрархічну структуру показників, що дозволяє звести завдання багатокритеріального оцінювання альтернатив до одного критерію або використовувати для вибору вектор показників;
- однозначність отриманої оцінки стану об'єкту;

- універсальність застосування за рахунок адаптації системи показників в ході роботи;
- не накопичує помилку навчання за рахунок використання процедур навчання;
- врахування типу невизначеності та зашумленості вихідних даних при побудові нечіткої когнітивної моделі;
- висока достовірність отриманих рішень при пошуку рішення у декількох напрямках;
- можливість об'єднання різнорідних вирішувачів без суворої прив'язки до композиції, що агрегує їх виходи;
- можливості вибору найкращої схеми комбінування вирішувачів.

До недоліків запропонованого методу слід віднести:

- втрата інформативності при оцінюванні стану об'єкту моніторингу за рахунок побудови функції належності;
- менша точність оцінювання по окремо взятому параметру оцінки стану об'єкту;
- втрата достовірності отриманих рішень при пошуку рішення в декількох напрямках одночасно;
- менша точність оцінювання у порівнянні з іншими методами оцінки.

Зазначений метод дозволить:

- провести оцінку стану об'єкту;
- визначити ефективні заходи для підвищення ефективності управління;
- підвищити швидкість оцінки стану об'єкту;
- зменшити використання обчислювальних ресурсів систем підтримки та прийняття рішень.

Запропонований підхід доцільно використовувати для вирішення задач оцінки складних та динамічних процесів, що характеризуються високим ступенем складності.

Зазначене дослідження є подальшим розвитком досліджень, що спрямовані на розробку методологічних засад підвищення ефективності інформаційно-

аналітичного забезпечення, що опубліковані вже раніше [345, 347–349, 366].

Напрямки подальших досліджень слід спрямувати на зменшення обчислювальних витрат при обробці різнотипних даних в системах спеціального призначення.

Висновки

1. Проведено формалізований опис задачі аналізу стану об'єктів в інформаційних системах спеціального призначення. Зазначена формалізація дозволяє описати процеси, що проходять в інформаційних системах спеціального призначення під час вирішення завдань аналізу стану об'єктів. В якості критерію ефективності зазначеного методу обрано оперативність процесу аналізу стану об'єкту при заданій достовірності. В ході дослідження сформульована концепція представлення методу оцінки в інформаційних системах спеціального призначення. В зазначеній концепції процес аналізу представлено у вигляді ієрархічного графу. Це дозволяє створити ієрархічний опис складного процесу за рівнями узагальнення та провести відповідний аналіз його стану.

2. У цьому дослідженні проведено розробку математичної моделі управління радіоресурсом систем радіозв'язку спеціального призначення на основі еволюційного підходу.

Результати дослідження стануть у нагоді при:

- розробці нових алгоритмів управління радіоресурсом;
- обґрунтуванні рекомендацій щодо підвищення ефективності оперативного управління радіоресурсом;
- аналізі радіоелектронної обстановки в ході ведення бойових дій (операцій);
- при створенні перспективних технологій підвищення ефективності оперативного управління радіоресурсом;
- оцінці адекватності, достовірності, чутливості науково-методичного апарату оперативного управління радіоресурсом;

– розробці нових та удосконаленні існуючих моделей управління радіоресурсом.

Напрямки подальших досліджень будуть спрямовані на розробку методології інтелектуального управління радіоресурсом систем радіозв'язку спеціального призначення.

3. Визначено алгоритм реалізації методу, що дозволяє:

враховується тип невизначеності та зашумленості даних;

врахувати наявні обчислювальні ресурси системи аналізу стану об'єкту;

провести точне навчання детекторів за рахунок комбінування процедур навчання (лінійне навчання та процедура навчання що розроблена в роботі [345]);

вибірковим задіяння ресурсів системи за рахунок підключення тільки необхідних типів детекторів;

побудувати класифікатор верхнього рівня за допомогою різних низькорівневих схем їх комбінування та агрегуючих композицій.

4. Проведений приклад використання запропонованої методу на прикладі оцінки стану оперативної обстановки угруповання війсь (сил). Зазначений приклад показав підвищення ефективності оперативності обробки даних на рівні 12–20 % за рахунок використання додаткових удосконалених процедур.