
CYBERSECURITY AND CRITICAL INFRASTRUCTURE PROTECTION

DOI 10.20535/2411-1031.2026.14.1.365500

УДК 004.056

ІГОР СВОБОДА,
ДМИТРО ЛАНДЕ**ПРІОРИТИЗАЦІЯ ЗАСОБІВ КОНТРОЛЮ БЕЗПЕКИ ДЛЯ КРИТИЧНОЇ ІНФРАСТРУКТУРИ ЗА ДОПОМОГОЮ МЕТОДУ АНАЛІЗУ МЕРЕЖ І ВЕЛИКИХ МОВНИХ МОДЕЛЕЙ**

У роботі вирішено актуальне науково-прикладне завдання пріоритизації засобів контролю безпеки об'єктів національної критичної інфраструктури (КІ) в умовах протидії загрозам з боку державних супротивників. Як методологічну основу дослідження обрано метод аналізу мереж (МAM), застосування якого дозволило змодельовати нелінійні взаємозалежності та цикли зворотного зв'язку в системі “засоби контролю – критерії оцінювання – вектори загроз – ресурсні обмеження”. Для реалізації поставленого завдання побудовано комплексну 20-вузлову модель, що охоплює чотири ключові кластери: засобів контролю (7 вузлів), критеріїв оцінювання (5 вузлів), векторів загроз (5 вузлів) та ресурсних обмежень (3 вузли). Науковою новизною роботи є запропонований підхід із використанням “віртуальних експертів”: необхідні для синтезу пріоритетів парні судження отримано від семи рольових персон, згенерованих великими мовними моделями (LLM). Ці персони імітують фахові погляди ключових стейкхолдерів – від інженера АСУ ТП та керівника SOC до CISO. Агрегація отриманих суджень методом середнього геометричного продемонструвала високу узгодженість експертних оцінок: середній Коефіцієнт Узгодженості (КУ) склав 0,006 ($\sigma = 0.004$), середній індекс Кочкодая – 0,034. За результатами побудови граничної суперматриці визначено глобальні пріоритети засобів безпеки. Лідируючі позиції посіли “Моніторинг” (0,1948) та “Сегментація” (0,1832), за якими з незначним відривом слідують “Ідентифікація та управління доступом” (РАМ) (0,1623), “Захист ланцюга постачання” (0,1354) та “Реагування на інциденти” (0,1342). Найменший пріоритет у рамках побудованої моделі отримало “Фізичне зміцнення” (0,0659). Стійкість отриманого ранжування підтверджено комплексним аналізом чутливості. Зокрема, метод LOEO засвідчив інваріантність двох найвищих пріоритетів (зміни рангів не перевищували $\pm 0,006$, коефіцієнт конкордації Кендалла W склав 0,77 $p < 0,001$), а моделювання Монте-Карло (1000 випробувань із Гауссовим шумом) довело надійність результатів. Практичну цінність запропонованої методології продемонстровано на прикладі портфельної моделі: при заданому бюджеті 10,2 млн дол. алгоритм сформував оптимальний набір засобів {Моніторинг, Ідентифікація, Реагування, Ланцюг постачання}, що забезпечує найкраще покриття у розрізі “засіб контролю – загроза”.

Ключові слова: метод аналізу мереж, кібербезпека, критична інфраструктура, державні загрози, багатокритеріальне прийняття рішень, великі мовні моделі, віртуальні експерти.

Постановка проблеми. Пріоритизація засобів контролю кібербезпеки для об'єктів національної критичної інфраструктури (КІ) характеризується винятковою складністю прийняття рішень. Оператори КІ постають перед необхідністю розподіляти обмежені ресурси – в умовах жорстких бюджетних та кадрових лімітів – таким чином, щоб максимально ефективно знизити ризики виконання завдань у перспективі 12–18 місяців, особливо в контексті протидії досвідченим зовнішнім державним супротивникам.

Аналіз існуючого методичного інструментарію свідчить, що традиційні методи пріоритизації, такі як метод аналізу ієрархій (MAI), обмежені своєю жорсткою, низхідною деревоподібною структурою. Такий підхід не дозволяє адекватно відобразити реалії сучасного середовища КІ, яке являє собою складну систему нелінійних взаємозалежностей та зворотних зв'язків. Наприклад, на ефективність засобу контролю “Моніторинг” суттєвий вплив має засіб “Сегментація”, а вартість впровадження нового засобу контролю може, у свою чергу, безпосередньо впливати на наявні бюджетні обмеження.

Окреслена проблематика зумовлює необхідність застосування методології, здатної моделювати ці складні цикли зворотного зв'язку. Метод аналізу мереж (МAM) [1], [2], який є узагальненням MAI, розроблений саме для цієї мети, оскільки його математичний апарат допускає наявність взаємозалежностей та зворотного зв'язку між кластерами елементів. Однак практичне застосування МAM стикається із суттєвим обмеженням: складністю отримання великої кількості узгоджених парних суджень від різноманітної групи експертів.

У даній роботі запропоновано вирішення обох зазначених проблем. Розроблено модель МAM, адаптовану спеціально для сфери безпеки КІ та АСУ ТП (Автоматизованої Системи Управління Технологічними Процесами), із фокусом на міжсекторальних практиках, що довели свою ефективність у зниженні ризиків (зокрема, сегментація, моніторинг, контроль облікових даних). Моделювання загроз у дослідженні спирається на матрицю та філософію MITRE ATT&CK for ICS. Для подолання проблем зі збором експертних суджень запроваджено новий метод, що використовує сім “віртуальних експертів” – рольових персон, згенерованих Великими Мовними Моделями (LLM). Вони імітують судження ключових зацікавлених сторін, таких як інженер АСУ ТП, директор з кібербезпеки чи керівник SOC. Запропонований підхід дозволяє швидко та прозоро генерувати повний набір парних порівнянь, необхідних для наповнення моделі, що забезпечує стійку, керовану даними пріоритизацію засобів контролю безпеки.

Аналіз останніх досліджень і публікацій. Теоретико-методологічне підґрунтя роботи сформовано на основі двох пов'язаних масивів літератури: фундаментальних праць з методології МAM, а також вітчизняних досліджень з мережевого аналізу та агрегації даних у кібербезпеці [3], [4]. Питання моделювання та оцінювання кіберзахисту і кіберстійкості об'єктів критичної (інформаційної) інфраструктури України також детально розглядаються у сучасних вітчизняних публікаціях [5], [6], [7]. Сучасний систематичний аналіз підходів до управління ризиками ІБ підкреслює практичну доцільність явного врахування взаємозалежностей між сімействами контролів, зокрема через поєднання методів причинно-наслідкового впливу та МAM-подібних механізмів рейтингування [8]. Крім того, відомі прикладні дослідження, що використовують багатокритеріальні методи прийняття рішень із залученням лінгвістичних суджень експертів для оцінювання рівня інформаційної безпеки промислових систем (ICS) [9].

Розглядаючи методологічну літературу з МAM, слід зазначити, що метод, розроблений Т. Л. Сааті, надає необхідну теоретичну основу для прийняття рішень в умовах залежностей. На відміну від MAI, МAM використовує суперматрицю для відображення впливів між кластерами та вузлами. Обчислювальний процес включає три етапи:

- формування незваженої суперматриці (побудованої з локальних власних векторів);
- розрахунок зваженої суперматриці (яка є стовпцево-стохастичною та масштабованою за впливами кластерів);
- отримання граничної суперматриці шляхом піднесення зваженої матриці до високого ступеня до досягнення збіжності.

Стовпці граничної матриці збігаються до головного власного вектора, що визначає глобальні пріоритети.

Критичним компонентом коректного застосування МАМ є валідація вхідних і вихідних даних. Матриці парних порівнянь підлягають обов'язковій перевірці на узгодженість, для чого використовується не лише стандартний Коефіцієнт Узгодженості (КУ) Сааті, але й індекс Кочкодая (K). Крім того, ретельній перевірці підлягає стійкість остаточних рангів. Аналіз чутливості є стандартним інструментом у методах багатокритеріального прийняття рішень (MCDM) для перевірки стійкості ранжування до варіацій ваг і похибок у судженнях [10]. У дослідженні застосовано багатоаспектну валідацію, що включає аналіз методом “ковзного ножа” (jackknife) за типом “Виключення по одному експерту” (LOEO) [11] та моделювання Монте-Карло зі збуренням початкових значень [12], що забезпечує надійність отриманих результатів.

Галузева література та настанови з безпеки КІ/АСУ ТП стали основою для формування набору засобів контролю та векторів загроз у моделі. Фундаментальним джерелом визначено “Посібник з безпеки Операційних Технологій” (NIST SP 800-82r3) [13], який визначає сегментацію, однонаправлені шлюзи та захист периметра як ключові механізми ізоляції, а також підкреслює критичну необхідність моніторингу. Зазначені положення доповнюються рамковим документом NIST CSF 2.0 [14], що забезпечує методологію пріоритизації, орієнтовану на ризики та кінцевий результат. Крім того, модель враховує “Міжсекторальні цільові показники кібербезпеки” (CPGs) від CISA [15], які надають базовий набір пріоритетних практик для всіх секторів КІ. Отриманий у роботі список засобів контролю узгоджується з цими показниками. Для обґрунтування моделі реалістичним ландшафтом загроз використано базу знань MITRE ATT&CK for ICS [16]. Вона каталогізує специфічні тактики, техніки та процедури (TTPs) державних суб'єктів (наприклад, початковий доступ, бічне переміщення), і саме здатність засобу контролю протидіяти цим конкретним TTPs визначає критерій “Ефективність” у запропонованій моделі МАМ.

Мета та завдання статті. Основною метою роботи є розробка та практична демонстрація науково обґрунтованої, строгої та відтворюваної методології пріоритизації інвестицій у засоби контролю кібербезпеки для операторів критичної інфраструктури, що функціонують в умовах протидії державним загрозам.

Для досягнення поставленої мети в роботі вирішено комплекс взаємопов'язаних завдань. По-перше, здійснено побудову комплексної мережевої моделі на основі методу аналізу мереж (МАМ), яка відображає системні взаємозалежності середовища безпеки КІ/АСУ ТП та включає визначення кластерів засобів контролю, векторів загроз, критеріїв оцінювання та ресурсних обмежень. По-друге, реалізовано процедуру отримання та агрегації парних суджень для всіх вузлів мережі із залученням “віртуальних експертів” на базі LLM, що дозволило імітувати різноманітні точки зору зацікавлених сторін. По-третє, проведено обчислення глобальних пріоритетів засобів контролю безпеки шляхом послідовного формування незваженої, зваженої та граничної суперматриць. По-четверте, виконано багатоаспектний аналіз чутливості та стійкості отриманого вектора пріоритетів із застосуванням методів LOEO, коефіцієнта конкордації W Кендалла та моделювання Монте-Карло. По-п'яте, здійснено трансформацію остаточного ранжування у план конкретних дій та побудовано практичну модель підтримки прийняття рішень, що враховує бюджетний портфель і сценарний аналіз.

Виклад основного матеріалу дослідження

Етап 1. Метод аналізу мереж (МАМ). Першим етапом дослідження є застосування методу аналізу мереж (МАМ). Зазначений метод узагальнює класичний метод аналізу ієрархій (МАІ) шляхом заміни його жорсткої ієрархічної будови на більш гнучку мережеву структуру. Такий підхід дозволяє ефективно моделювати складні взаємозалежності та цикли зворотного зв'язку, притаманні досліджуваній предметній області.

У межах моделі МАМ вузли групуються у функціональні кластери, а ступінь впливу одного вузла на інший кількісно оцінюється за допомогою процедури парних порівнянь. Отримані локальні пріоритети слугують основою для формування незваженої суперматриці

W . Надалі ця матриця зважується з використанням матриці міжкластерного впливу C для створення стохастичної за стовпцями зваженої суперматриці $W^{(w)}$.

Шляхом піднесення отриманої матриці до достатньо високого степеня ($k \rightarrow \infty$) досягається її збіжність до граничної суперматриці $W^{(\infty)}$, стовпці якої стають ідентичними. При цьому кожен стовпець відповідає стаціонарному розподілу π , який задовольняє умову $W^w \pi = \pi$ (що представляє собою правий власний вектор для власного значення 1). Остаточні глобальні пріоритети альтернатив (засобів) контролю визначаються на основі елементів вектора π , що відповідають вузлам кластера “Засоби Контролю”. Варто зазначити, що у випадках, коли матриця є звідною або періодичною, для розрахунків використовується границя Чезаро (Cesàro limit) [17].

Етап 2. Мережева модель. Розроблена в рамках дослідження мережева модель охоплює 20 вузлів, які згруповано у чотири функціональні кластери. Перший кластер, “Засоби Контролю”, містить сім альтернатив, що підлягають пріоритизації:

- 1) моніторинг мережі АСУ ТП та виявлення аномалій (MON);
- 2) сегментація мережі та однонаправлені шлюзи (SEG);
- 3) ідентифікація, MFA та управління привілейованим доступом (IDM);
- 4) реагування на інциденти та навчання персоналу (IR);
- 5) забезпечення безпеки ланцюга постачання, включно з SBOM та підписанням коду (SUP) [18];
- 6) автономні незмінні резервні копії та відновлення для АСУ ТП (BAK);
- 7) фізичне та середовищне зміцнення об’єктів (PHY).

Другий кластер, “Критерії”, визначає п’ять метрик для оцінювання засобів контролю, а саме:

- 1) ефективність проти тактик, технік і процедур (TTPs) державних суб’єктів (E);
- 2) час на впровадження (T);
- 3) вплив на операційну діяльність та перешкоди доступності (U);
- 4) вартість реалізації (C);
- 5) дотримання політик та комплаєнсу (P), наприклад, узгодженість із CSF/NIS2.

Третій кластер, “Вектори Загроз”, охоплює п’ять ключових джерел ризику:

- 1) цільовий спінфішинг (PHISH);
- 2) компрометація ланцюга постачання (SC);
- 3) дистанційна експлуатація вразливостей (EXP);
- 4) внутрішня загроза або інсайдер (INS);
- 5) фізичний саботаж (SAB).

Четвертий кластер, “Обмеження”, відображає реальні перешкоди впровадження:

- 1) бюджетні обмеження (BUD);
- 2) готовність персоналу (WRK);
- 3) ризик прив’язки до конкретного постачальника (LOCK).

Матриця “Кластер-Кластер” (C). Для формалізації взаємовпливів між чотирма джерельними кластерами (Засоби Контролю, Загрози, Критерії, Обмеження) отримано глобальний апіорний вектор впливу $v = [0,4466, 0,2549, 0,2506, 0,0480]^T$. Зазначений вектор відображає специфіку предметної області, де впливи, орієнтовані на засоби контролю, домінують у процесі прийняття рішень, вторинний внесок мають дані про загрози та критерії, а роль обмежень є меншою, проте значущою.

На основі вектора v для кожного цільового кластера j формується стовпець $C_{.j}$. Процедура формування передбачає обнуління елементів у випадках, коли відповідний блок суперматриці $W_{.j}$ є структурно відсутнім, та ренормування решти елементів до одиниці за формулою:

$$C_{.j} = \text{renorm}(v \odot 1\{W_{.j} \neq 0\}). \quad (1)$$

Такий підхід дозволяє отримати специфічні для стовпця ваги кластерів, які:

- враховують структурні нулі;
- гарантують, що кожен блоковий стовпець зваженої суперматриці $W^{(w)}$ залишається стохастичним;
- дозволяють уникнути введення вільних параметрів, не підтверджених емпіричними даними.

Розраховані ваги наведено в Таблиці 1. Емпіричний аналіз підтвердив стійкість моделі до варіацій вектора v : лідируючі позиції (Моніторинг, Сегментація) залишаються стабільними при помірних пропорційних збуреннях, а незначні зміни рангів у середині списку узгоджуються з варіативністю методу LOEO (деталізовано на етапі 6).

Таблиця 1 – Глобальний апіорний вектор впливу кластерів v , що використовується для параметризації ваг, специфічних для стовпців

Вплив ↓ / Під впливом →	Засоби Контролю	Загрози	Критерії	Обмеження
Засоби Контролю	0,4466	0,4466	0,4466	0,4466
Загрози	0,2549	0,2549	0,2549	0,2549
Критерії	0,2506	0,2506	0,2506	0,2506
Обмеження	0,0480	0,0480	0,0480	0,0480

Структура впливу. Побудована модель явно визначає цикли зворотного зв'язку та системні взаємозалежності. Зокрема, прямі впливи включають напрямки:

- Засоби Контролю → Критерії;
- Загрози → Критерії;
- Обмеження → Критерії та → Засоби Контролю.

Зворотний зв'язок моделюється через внутрішню залежність у кластері Критерії ↔ Критерії, а також вплив Засобів Контролю на Обмеження. Окрім цього, модель враховує ключові синергії типу “Засіб Контролю ↔ Засіб Контролю”, такі як:

- SEG → MON (покращення спостережності);
- IDM → MON (збагачення контексту виявлення);
- SEG → IDM (зменшення поверхні атаки на системи РАМ);
- ВАК → IR (підвищення ефективності реагування).

У випадках відсутності впливу відповідний блок W_{ij} приймається рівним нулю.

Етап 3. Методологія отримання суджень та обчислень. Отримання суджень на основі персон. Для наповнення моделі даними застосовано метод “віртуальних експертів”. Створено сім рольових персон на базі великих мовних моделей (LLM), які надавали парні судження за шкалою відношень Сааті (1-9). Детальний опис персон та інструкцій для їхнього обумовлення (prompts) наведено в Таблиці 2. Агрегація індивідуальних суджень здійснювалася за допомогою середнього геометричного, що забезпечує збереження властивості зворотності матриць.

Варто зазначити, що такий підхід узгоджується із сучасними дослідженнями. Зокрема, доведено можливість використання LLM для еліцитації формалізованих експертних знань і побудови апіорних припущень у кількісних моделях [19]. Крім того, існують підходи до масштабування групового прийняття рішень, де LLM застосовуються для кластеризації та узгодження переваг учасників [20], що концептуально підтримує обрану методологію.

Таблиця 2 — Обумовлення персон віртуальних експертів

ID	Роль (скорочено)	Обумовлення / Інструкції
E1	Інженер АСУ ТП	“Ви експлуатуєте ПЛК/РСК (PLC/DCS) різних виробників... Пріоритизуйте безпеку/доступність; не довіряйте інтрузивному скануванню. ...менше часу/витрат/впливу = краще”.
E2	Керівник SOC (blue team)	“Ви керуєте цілодобовим SOC з датчиками DPI в АСУ ТП... Цінують затримку виявлення та навантаження на сортування (triage)”.
E3	Спеціаліст з 'red team' АСУ	“Емуляція супротивника; зважуйте початковий доступ, бічне переміщення... враховуйте діоди даних та проміжні хости (jump hosts)”.
E4	Керівник з ланцюгів постачання	“Забезпечення SBOM/підписання коду; ризики третіх сторін; походження оновлень (provenance)”.
E5	Спеціаліст з комплаєнсу	“Зіставлення з функціями NIST CSF 2.0 та результатами NIS2; зважуйте відповідність політиці Р”.
E6	Менеджер польових операцій	“Технічні спеціалісти, дозволи на роботи, вікна обслуговування; зважуйте операційні перешкоди (friction) U”.
E7	Директор з кібербезпеки	“Бюджет BUD, персонал WRK, прив'язка до постачальника LOCK; зважуйте управління та послідовність (впровадження)”.

Відтворюваність та математична модель. З метою забезпечення стабільності та відтворюваності результатів усі запити до LLM виконувалися з фіксованими параметрами декодування (температура = 0,2, $top_p = 1,0$) та вимогою структурованого числового виходу. Для генерації суджень віртуальних експертів використовувалася модель gpt-5-2025-08-07 від OpenAI через API станом на 1 жовтня 2025 року.

Розрахунок локальних пріоритетів та узгодженості. Для кожної зворотної матриці парних порівнянь A (де $a_{ij} = 1/a_{ji}$) локальний вектор пріоритетів w визначався як нормований головний правий власний вектор. Він знаходиться шляхом розв'язання рівняння $Aw = \lambda_{\max} w$ за умови нормування суми елементів до одиниці [2]. Хоча існують методи роботи з неповними матрицями [21], у даному дослідженні використано повні матриці, сформовані віртуальними експертами.

Валідація отриманих суджень включала розрахунок Індексу Узгодженості (CI) за формулою $CI = (\lambda_{\max} - n) / (n - 1)$ та Коефіцієнта Узгодженості (CR) $CR = CI / RI_n$ [2], а також обчислення індексу Кочкодая $K(A)$ [22]. Усі матриці задовольнили порогові значення $CR < 0,10$ та $K < 0,10$. Середній показник CR для отриманих матриць склав 0,006, що свідчить про високу узгодженість. Для критеріїв вартості, часу та впливу (T, C, U) застосовано шкалу корисності (“менше – краще”). Отримані вектори пріоритетів для всіх п'яти критеріїв наведено в Таблиці 3.

Зважування суперматриці. На основі локальних векторів сформовано незважену суперматрицю W , де кожен блок W_{ij} відображає вплив вузлів кластера i на вузли кластера j . Кожен ненульовий підстовпець у W_{ij} сумується до 1. Зважена суперматриця $W^{(w)}$ отримана шляхом блокового масштабування з використанням матриці міжкластерного впливу C :

$$(W^{(w)})_{ij} = C_{ij} \cdot W_{ij}. \quad (2)$$

Для забезпечення стохастичності матриці за стовпцями реалізовано спеціальну процедуру обробки структурних нулів. Для будь-якого цільового кластера j матриця C налаштовується так, що $C_{ij} = 0$, якщо $W_{ij} = 0$ а сума елементів стовпця $C_{\cdot j}$ по активній множині $A(j) = \{i : W_{ij} \neq 0\}$ дорівнює одиниці ($\sum_i C_{ij} = 1$). Це гарантує, що сума елементів кожного стовпця зваженої суперматриці дорівнює одиниці:

$$\sum_{i \in A(j)} C_{ij} \cdot (\text{сума стовпця } W_{ij}) = \sum_{i \in A(j)} C_{ij} \cdot 1 = 1. \quad (3)$$

Отримана матриця $W^{(w)}$ є стохастичною за стовпцями, що є необхідною умовою для коректного обчислення границі [1], [17], та концептуально наближає метод до алгоритмів типу PageRank [23].

Обчислення границі та діагностика ергодичності. Оскільки індукований граф моделі є сильно зв'язним та аперіодичним (завдяки наявності зворотних зв'язків та синергій), матриця $W^{(w)}$ є примітивною. Це означає, що послідовність $W^{(w)k}$ збігається до границі рангу один. Така поведінка підтверджується наявністю спектрального проміжку ($|\lambda_2| < 1$).

Граничну суперматрицю $W^{(\infty)}$ обчислено методом повторних піднесення до квадрата до досягнення умови збіжності $\|W^{(w)2^k} - W^{(w)2^{k-1}}\|_1 < 1 \times 10^{-10}$. Стаціонарний розподіл π отримано з довільного стовпця матриці $W^{(\infty)}$. Числова верифікація шляхом прямого розв'язання задачі на власний вектор $W^w \pi = \pi$ з підтвердила результат з точністю 4.1×10^{-12} . Додаткова перевірка з використанням демпфірування $\delta = 10^{-6}$ показала ідентичний результат з точністю 10^{-6} .

Етап 4. Вибіркові локальні результати. Обчислення кінцевої моделі базувалося на використанні локальних векторів пріоритетів, отриманих в результаті агрегації суджень віртуальних експертів. Варто зазначити, що використані дані містили незначні, реалістичні неузгодженості, що відрізняє їх від теоретично “ідеальних” матриць. Ключові локальні вектори, які слугували основою для наповнення суперматриці, систематизовано в Таблиці 3.

Необхідно окремо зауважити, що рядок “Критерії (Е, Т, U, С, Р)” у наведеній таблиці відображає розподіл ваг усередині самого кластера “Критерії” і не має функціонального зв'язку з апіорним вектором кластера v , визначеним на другому етапі дослідження. Спостережувана числова близькість між значенням ефективності $E = 0,4463$ та вагою кластера засобів контролю $v_{zc} = 0,4466$ є виключно випадковим збігом.

Таблиця 3 – Вибіркові локальні вектори пріоритетів (з агрегованих експертних суджень)

Критерій	Засіб Контролю	w_{local}	λ_{max}	CI
Ціль	Критерії (Е, Т, U, С, Р)	[0,4463, 0,1504, 0,1016, 0,1475, 0,1543]	5,0030	0,00067
Ефективність (Е)	Засоби Контролю	[0,2479, 0,2737, 0,1299, 0,1710, 0,0972, 0,0506, 0,0297]	7,0123	0,00155
Час (Т)	Засоби Контролю	[0,0962, 0,0693, 0,1438, 0,1804, 0,1973, 0,2126, 0,1004]	7,0090	0,00114
Вплив (U)	Засоби Контролю	[0,0920, 0,1006, 0,1618, 0,1743, 0,1710, 0,2016, 0,0987]	7,0100	0,00126
Вартість (С)	Засоби Контролю	[0,0864, 0,0716, 0,1338, 0,1772, 0,1712, 0,2359, 0,1239]	7,0143	0,00180
Політика (Р)	Засоби Контролю	[0,1564, 0,1817, 0,1696, 0,1291, 0,2095, 0,0936, 0,0600]	7,0067	0,00084

Етап 5. Основні результати: Глобальні пріоритети. Граничний вектор, що визначає глобальні пріоритети, було обчислено для стохастичної за стовпцями матриці $W^{(w)}$. Враховуючи, що значення другого власного числа $|\lambda_2| \approx 0,947$ вказує на повільну збіжність при простій ітерації, для досягнення швидкої чисельної стабільності було застосовано метод повторних піднесення до квадрата. Ітераційний процес тривав до моменту, поки зміна за L_1 -нормою не стала меншою за порогове значення 1×10^{-10} .

Для підтвердження надійності результатів проведено перевірку на стійкість шляхом формування демпфированої матриці $\tilde{W} = (1 - \delta)W^{(w)} + \delta \cdot 11^T / N$ з параметром $\delta = 10^{-6}$. Отриманий стаціонарний вектор відрізнявся від недемпфированої границі менш ніж на 10^{-6} у кожній координаті. Додаткова перехресна перевірка через розв'язання рівняння на правий власний вектор $W^w \pi = \pi$ (за умови нормування $\sum_i \pi_i = 1$) продемонструвала збіг з границею з високою точністю 4.1×10^{-12} (за L_∞ -нормою).

Кінцеві глобальні пріоритети для семи досліджуваних засобів контролю безпеки, отримані на основі граничної суперматриці $W^{(\infty)}$, представлені в Таблиці 4 та візуалізовані на Рисунку 1.

Таблиця 4 – Кінцеві глобальні пріоритети для засобів контролю безпеки

Ранг	Засіб контролю (абрев.)	Глобальна вага
1	Моніторинг мережі АСУ ТП та виявлення аномалій (MON)	0,1948
2	Сегментація мережі та однонаправлені шлюзи (SEG)	0,1832
3	Ідентифікація та управління привілейованим доступом (IDM)	0,1623
4	Забезпечення безпеки ланцюга постачання / SBOM / підписання коду (SUP)	0,1354
5	Реагування на інциденти та навчання (IR)	0,1342
6	Автономні / незмінні резервні копії та відновлення АСУ ТП (BAK)	0,1242
7	Фізичне / середовищне зміцнення (PHY)	0,0659

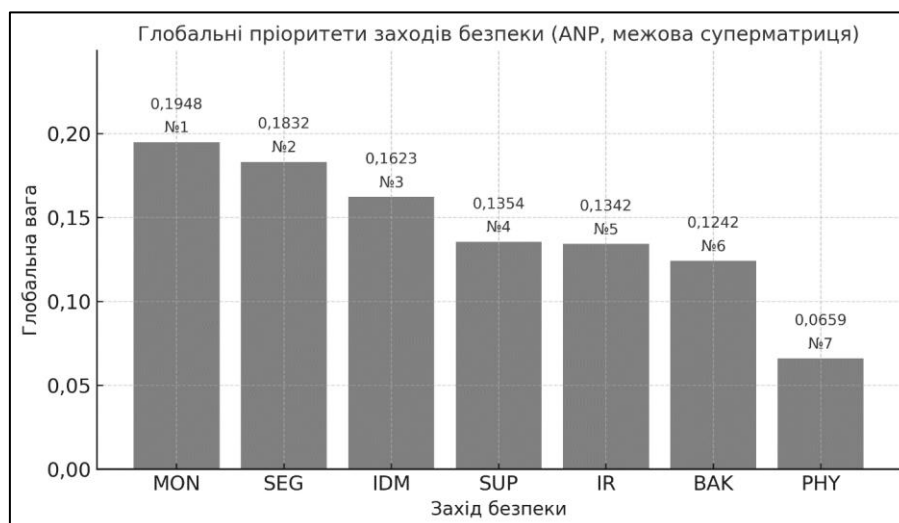


Рисунок 1 – Глобальні пріоритети засобів контролю безпеки (гранична суперматриця MAM)

Інтерпретація результатів. Отримані дані дозволяють виділити чітку пріоритетну групу засобів контролю. Найвищі позиції посідають “Моніторинг” та “Сегментація”, що повністю узгоджується з офіційними міжнародними настановами, які акцентують увагу на критичності ізоляції та виявлення загроз у середовищах АСУ ТП (ОТ) [13]-[16]. Наступними за важливістю є засоби ідентифікації (РАМ), а також пара засобів із близькими рангами – безпека ланцюга постачання та реагування на інциденти. Замикає рейтинг фізичне зміцнення, яке, залишаючись необхідним елементом захисту, має найнижчий пріоритет з точки зору розподілу інвестицій у рамках даної моделі.

Етап 6. Аналіз чутливості та стійкості. З метою підтвердження достовірності та обґрунтованості отриманого ранжування в роботі виконано комплексний аналіз стійкості, що включав три рівні перевірки.

Метод виключення по одному експерту (LOEO). На першому етапі застосовано метод LOEO. Процедура передбачала перерахунок моделі сім разів, причому на кожній ітерації вилучалися судження одного з віртуальних експертів. Отримані результати, узагальнені в Таблиці 5, демонструють виняткову стійкість моделі. Зокрема, два провідні засоби контролю (MON, SEG) залишаються інваріантними в усіх ітераціях, а IDM стабільно утримує третю позицію. Незначні перестановки в середині рейтингу спостерігаються лише між SUP та IR, проте коливання їхніх глобальних ваг відбуваються у дуже вузькому діапазоні ($SD \approx 0,006$). Розрахункова середня абсолютна зміна глобальних ваг становить 0,0039, а середній коефіцієнт τ Кендалла відносно базового ранжування дорівнює 0,90.

Графічна інтерпретація стійкості LOEO наведена на Рисунку 2, де маркери (квадрати) позначають базові ваги, кола – середні значення за ітераціями LOEO ($\pm$ середньоквадратичне відхилення), а горизонтальні відрізки візуалізують діапазон min-max. Варто зазначити, що оскільки всі персони згенеровані на базі одного сімейства LLM, цей метод оцінює чутливість саме до формулювання запитів та рольового обумовлення. Показник “Зміни рангу” в Таблиці 5 відображає кількість ітерацій, у яких позиція засобу відрізнялася від базової.

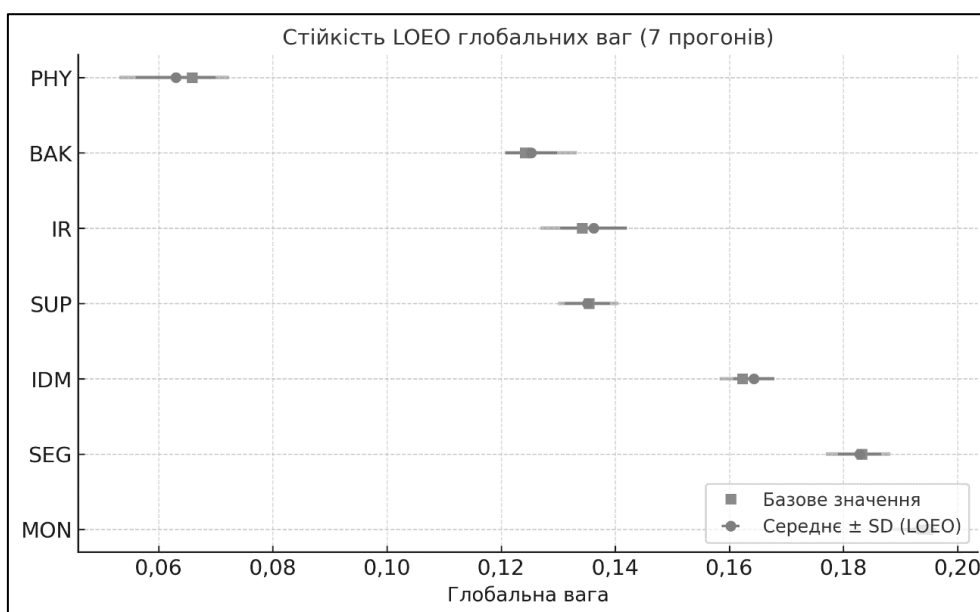


Рисунок 2 – Стійкість LOEO для глобальних ваг (7 ітерацій)

Міжекспертна узгодженість. На другому етапі оцінено рівень узгодженості між сімома віртуальними експертами (оцінювачами, $m = 7$) щодо ранжування $N = 7$ засобів контролю. Для цього використано коефіцієнт конкордації Кендалла W [24]. Розрахунки показали, що для всіх критеріїв узгодженість є суттєвою та статистично значущою ($p < 0,001$, визначено за

асимптотичною χ^2 -апроксимацією). Зокрема, отримано такі значення коефіцієнта: ефективність E ($W=0,77$), P ($W=0,74$), вартість C ($W=0,68$), час T ($W=0,65$) та операційний вплив U ($W=0,62$).

Таблиця 5 – Аналіз чутливості LOEO для глобальних ваг

Засіб контролю	Базовий рівень	Середнє LOEO $\pm$ СКВ	Min-Max (7 ітерацій)	Зміни рангу (з 7)
MON	0,1948	0,1933 $\pm$ 0,0028	0,1887–0,1964	0
SEG	0,1832	0,1828 $\pm$ 0,0038	0,1769–0,1882	0
IDM	0,1623	0,1643 $\pm$ 0,0036	0,1583–0,1679	0
SUP	0,1354	0,1351 $\pm$ 0,0040	0,1300–0,1406	4
IR	0,1342	0,1362 $\pm$ 0,0058	0,1269–0,1421	5
BAK	0,1242	0,1253 $\pm$ 0,0045	0,1207–0,1333	1
PHY	0,0659	0,0630 $\pm$ 0,0070	0,0531–0,0723	0

Збурення методом Монте-Карло. Третім етапом перевірки стало стохастичне моделювання методом Монте-Карло, в рамках якого проведено 1000 випробувань в області відношень. Сутність методу полягала у внесенні збурень у кожне судження за формулою $\log a_{ij} \leftarrow \log a_{ij} + \epsilon_{ij}$ де $\epsilon_{ij} \sim N(0, \sigma^2)$, із забезпеченням зворотності матриці $a_{ji} = 1/a_{ij}$. Дисперсія σ обиралася таким чином, щоб згенеровані матриці залишалися в межах типового для ретельної експертизи діапазону узгодженості ($CR \approx 0,05 - 0,10$). Для кожного випробування відбувалися повна перебудова зваженої суперматриці (з урахуванням специфічних матриць) та перерахунок граничного вектора. Результати моделювання підтвердили, що найвища група пріоритетів (MON, SEG) залишається інваріантною, а перестановки в середині списку є обмеженими та корелюють із результатами аналізу LOEO. Зокрема, у 1000 випробуваннях пара лідерів (MON, SEG) зберігала топ-2 позиції у 100% випадків, а перестановка між SUP та IR спостерігалася у 57% випадків. Максимальна зміна ваги для будь-якого засобу не перевищувала 0,015.

Етап 7. Практичне застосування та підтримка прийняття рішень. Практичний план дій. Отримане в результаті моделювання ранжування надає основу для формування детального плану впровадження заходів безпеки, який охоплює сім ключових напрямків.

По-перше, критичним є напрямок “Виявляти” (MON), що передбачає розгортання пасивних датчиків DPI (Deep Packet Inspection) безпосередньо в середовищі АСУ ТП, а також впровадження спеціалізованих систем виявлення аномалій.

По-друге, необхідно “Ізолювати” (SEG) мережеві ресурси, для чого слід забезпечити жорстку сегментацію на рівні технологічних клітинок або зон та встановити однонаправлені шлюзи (data diodes) на межах периметра.

По-третє, напрямок “Контролювати Ідентичності” (IDM) вимагає впровадження систем управління привілейованим доступом (PAM) для інженерного складу та обов’язкової багатofакторної автентифікації (MFA) для будь-якого дистанційного доступу.

По-четверте, слід “Захищати Ланцюг” (SUP) шляхом запровадження використання специфікацій програмних компонентів (SBOM) та суворої перевірки цифрових підписів коду. Варто зазначити, що емпіричне дослідження практик SBOM у програмній інженерії демонструє, що навіть за умови визнання SBOM як ключового артефакту прозорості ланцюга постачання, реальне впровадження часто стикається зі значними бар’єрами процесів, інструментарію та розподілу відповідальності сторін [25].

По-п'яте, необхідно “Готуватися до Реагування” (IR), розробивши специфічні для середовищ АСУ ТП інструкції з реагування на інциденти та регулярно проводячи командно-штабні навчання.

По-шосте, напрямок “Відновлювати” (BAK) передбачає створення незмінних (immutable) та автономних резервних копій для контролерів і людино-машинних інтерфейсів (HMI).

Нарешті, по-сьоме, слід “Зміцнювати Фізично” (PHY) об'єкти, встановивши замки на апаратні шафи та системи виявлення несанкціонованого фізичного доступу або втручання.

Модель підтримки прийняття рішень. З метою трансформації отриманих пріоритетів у економічно обґрунтовані програми, задачу було змодельовано як задачу про “рюкзак” із субмодулярним покриттям (submodular coverage knapsack). Для кожної загрози t , що характеризується очікуваними річними втратами EAL_t , та кожного засобу контролю c , що забезпечує часткове покриття $\alpha_{c,t}$, залишкові втрати після розгортання набору засобів $S \subseteq C$ визначаються формулою:

$$L(S) = \sum_t EAL_t \cdot \prod_{c \in S} (1 - \alpha_{c,t}). \quad (4)$$

Відповідно, функція зменшення ризику $f(S) = \sum_t EAL_t - L(S)$ є монотонною та субмодулярною. Максимізацію функції $f(S)$ в умовах бюджетних обмежень реалізовано з використанням “лінивого” жадібного (lazy greedy) алгоритму [26]. Відомо, що для максимізації монотонних субмодулярних функцій жадібний алгоритм досягає гарантії ефективності на рівні $(1 - 1/e)$ при обмеженні на потужність, а для бюджетів типу “рюкзак” (knapsack budgets) [27] доступні майже оптимальні гарантії з константним коефіцієнтом (constant factor guarantees) через варіанти з частковим перебором. У даній роботі реалізовано стандартну “ліниву” жадібну схему, яка демонструє очікувану поведінку зниження граничної віддачі [28].

Примітка щодо калібрування. Коефіцієнти покриття $\alpha_{c,t}$ були отримані на основі зіставлення “загроза-контроль” (з використанням бази знань ATT&CK for ICS) та експертного калібрування, тоді як вектор витрат k відображає однорічні витрати на програму, отримані з офіційної документації постачальників та типових розподілів бюджету КІ. Обидва параметри узгоджуються з пріоритетами, визначеними моделлю МАМ. Значення EAL є умовними і використовуються лише для демонстрації.

З вектором витрат $k = \{3, 8, 5, 2, 3, 0, 1, 4, 2, 0, 1, 6, 0, 8\}$ (млн дол.) та ідентичною матрицею покриття, евристичний алгоритм обирає набір {MON, IDM, IR, SUP} при бюджеті $B=10,2$ млн дол. Це забезпечує змодельоване зменшення ризику на ≈ 27 млн дол. Аналіз чутливості показав, що зміна коефіцієнтів $\alpha_{c,t}$ на $\pm 20\%$ змінює загальну вигоду менш ніж на 4% і зберігає ранжування незмінним.

Аналіз сценаріїв. Додатково було проведено перезважування моделі з метою оцінки зміщення пріоритетів у випадку, якщо певні загрози стануть домінуючими (результати наведено в Таблиці 6). Процедура перезважування для сценарію S полягала у посиленні визначених загроз шляхом масштабування блокового стовпця “Загрози $\rightarrow$ Критерії” за формулою:

$$C_{\cdot, \text{Критерії}} \leftarrow \text{renorm}(s \odot C_{\cdot, \text{Критерії}}), \quad (5)$$

де S – вектор додатних множників для п'яти вузлів Загроз (посилені загрози мають множник $s > 1$, інші – $s < 1$),

“renorm” – операція, що перемасштабує активні елементи так, щоб їх сума дорівнювала одиниці.

Усі інші блоки та стовпці в матриці C залишаються незмінними, а зважена суперматриця $W^{(w)}$ перебудовується з використанням тієї ж схеми зважування з урахуванням структурних нулів (етапи 2-3). Після цього граничний вектор перераховується на основі оновленої $W^{(w)}$.

Таблиця 6 – Аналіз сценаріїв для глобальних ваг

Сценарій	MON	SEG	IDM	SUP	IR	BAK	PHY
Базовий рівень	0,1948	0,1832	0,1623	0,1354	0,1342	0,1242	0,0659
Домінування SC	0,188	0,177	0,150	0,160	0,136	0,116	0,073
Домінування EXP	0,196	0,198	0,164	0,121	0,126	0,098	0,097
Домінування INS+SAB	0,182	0,169	0,176	0,110	0,160	0,117	0,086

Розглянемо приклад розрахунку для сценарію “Домінування загрози Ланцюгів Постачання (SC)”. У вихідній (базовій) моделі трьома основними пріоритетами були: Моніторинг (MON) – 0,1948, Сегментація (SEG) – 0,1832 та Ідентифікація (IDM) – 0,1623. Після переважування, де модель була перерахована зі збільшеними “вагами стовпця загроз” для вектора загрози “Ланцюг постачання” (імітація сценарію домінування цієї загрози), пріоритети змістилися. Нові результати склали:

- Моніторинг (MON) – 0,188;
- Сегментація (SEG) – 0,177;
- Ланцюг постачання (SUP) – 0,160.

Таким чином, у сценарії “Домінування SC” засіб контролю “Ланцюг постачання” (SUP), який займав 4-те місце в базовій моделі, піднявся і став 3-м за пріоритетом, витіснивши “Ідентифікацію” (IDM). Аналогічно, у сценарії домінування експлойтів (EXP) засіб SEG стає абсолютним лідером (№1). Це демонструє гнучкість моделі для планування з урахуванням динаміки загроз.

Висновки. У даній роботі успішно розроблено та апробовано стійку, прозору та багатоаспектну методологію пріоритизації засобів контролю кібербезпеки об’єктів критичної інфраструктури в умовах протидії державним загрозам. Завдяки синергетичному поєднанню методу аналізу мереж (МАН) із новітнім підходом до отримання суджень від «віртуальних експертів» на базі LLM, вдалося ефективно змодельовати складні нелінійні взаємозалежності середовища АСУ ТП та отримати стабільне, науково обґрунтоване ранжування.

Основним практичним результатом дослідження є чітка пріоритизація засобів контролю. Визначено, що Моніторинг мережі АСУ ТП (0,1948) та Сегментація мережі (0,1832) є критично важливими напрямками інвестування. Наступні позиції посідають Ідентифікація/РАМ (0,1623), Забезпечення безпеки ланцюга постачання (0,1354) та Реагування на інциденти (0,1342). Фізичне зміцнення (0,0659), залишаючись необхідним, отримало найменший відносний пріоритет. Виняткова стійкість отриманого ранжування підтверджена комплексним аналізом: метод LOEO засвідчив інваріантність двох лідируючих позицій, а моделювання Монте-Карло продемонструвало їх збереження навіть за умов стохастичних збурень. Вагомим підтвердженням валідності моделі є сильна кореляція результатів із галузевими настановами NIST, CISA та MITRE.

Наукова новизна та переваги запропонованої методології полягають у наступному. По-перше, на відміну від класичного МАІ, застосований підхід відображає системний зворотний зв’язок та взаємозалежності, що є критичним для адекватного моделювання кібербезпеки. По-друге, використання “віртуальних експертів” LLM забезпечує швидкість та відтворюваність збору суджень при збереженні високої узгодженості (середній $CR = 0,006$). По-третє, надійність результатів гарантується багатокомпонентним аналізом чутливості (LOEO, W Кендалла, Монте-Карло). По-четверте, методологія має пряме прикладне застосування в

інструментах підтримки прийняття рішень, зокрема для портфельної оптимізації та сценарного аналізу.

Обмеження та перспективи подальших досліджень

Обмеження. Необхідно відзначити певні обмеження проведеного дослідження. Аналіз LOEO було виконано на семи “віртуальних експертах”, згенерованих однією базовою великою мовною моделлю. Це означає, що перевірка стосувалася стабільності відносно формулювання запитів, проте не враховувала потенційні системні упередження, притаманні конкретному сімейству моделей. Висока узгодженість (наприклад, W Кендалла = 0,77 для критерію “Ефективність”) може частково пояснюватися цією однорідністю. Для підвищення зовнішньої валідності доцільно диверсифікувати пул експертів, залучаючи різні сімейства моделей, що заплановано у подальшій роботі.

Напрямки подальших досліджень. Результати роботи та виявлені обмеження окреслюють низку перспективних напрямків для майбутніх розвідок:

1. *Розширений аналіз чутливості.* Планується вдосконалення методу стійкого збурення. Оскільки поточне моделювання Монте-Карло працює в просторі логарифмічних відношень ($\log a_{ij} \leftarrow \log a_{ij} + \varepsilon$, де $\varepsilon \sim N(0, \sigma^2)$), корисним розширенням стане калібрування σ відповідно до типових профілів людської неузгодженості та аналіз емпіричних розподілів CR та K (індексу Кочкодая) при варіюванні σ (σ sweeps). Також доцільно провести глобальний аналіз чутливості, використовуючи діаграми “торнадо” для візуалізації впливу ваг “кластер-кластер” (C_{ij}) на кінцеві пріоритети.

2. *Підвищення різноманітності LOEO.* Аналіз “виключення по одному” слід масштабувати до схеми “виключення одного сімейства моделей” для виявлення фундаментальних упереджень. Альтернативним підходом може стати bootstrap-аналіз за варіантами запитів для перевірки лінгвістичної стійкості.

3. *Поглиблена звітність.* Майбутні дослідження мають включати не лише оцінку стабільності рангів, але й повні розподіли рангової кореляції (τ Кендалла) та кореляції ваг (ρ Спірмена) на основі випробувань зі збуренням. Окремим завданням є аналіз методом виключення демпфірування для оцінки залежності вектора пріоритетів від коефіцієнта δ (наприклад, для $\delta \in \{0, 10^{-8}, 10^{-6}, 10^{-4}, 10^{-2}\}$) [17]. Це дозволить кількісно оцінити L_1 -збурення граничного вектора та підтвердити коректність вибору $\delta = 10^{-6}$.

4. *Методологічне розширення.* Перспективним є масштабування моделі до повної структури BOCR-ANP (Вигоди, Можливості, Витрати, Ризики) з окремими суперматрицями. Крім того, критерій “Ефективність” слід посилити через створення керованої даними матриці покриття “Загроза – Засіб контролю”, що базується на зіставленні конкретних TTPs (ATT&CK for ICS) із засобами протидії. Валідацію рекомендацій моделі слід проводити на основі ретроспективного аналізу реальних інцидентів, таких як атаки TRITON або Colonial Pipeline [29], [30].

Внесок авторів:

– Ігор Свобода – методологія; аналіз джерел; проведення обчислювальних експериментів; аналіз результатів; підготовка первинного тексту статті; візуалізація; редагування та доопрацювання рукопису.

– Дмитро Ланде – наукове керівництво; концептуалізація; методологія; формальна постановка задачі; перевірка результатів. Усі автори прочитали та погодили остаточну версію рукопису.

Декларація щодо використання штучного інтелекту. Під час підготовки цієї роботи автори використовували велику мовну модель GPT-5 як інструмент дослідження для формування суджень віртуальних експертів і побудови парних порівнянь; відповідні процедури, промпти та результати описано в основному тексті статті. Також інструмент генеративного штучного інтелекту OpenAI Codex використовувався для технічної допомоги

під час підготовки рукопису: перевірки мовного оформлення, форматування, приведення формул і бібліографічних посилань до вимог журналу. Після використання цих інструментів автори перевірили та відредагували матеріал у необхідному обсязі й несуть повну відповідальність за зміст публікації. Наукові висновки, інтерпретація результатів і рекомендації сформульовані авторами.

Конфлікт інтересів. Автори заявляють про відсутність конфлікту інтересів. Під час підготовки цієї роботи не було комерційних, фінансових або інших відносин, які могли б вплинути на результати дослідження або їх інтерпретацію.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

- [1] T.L. Saaty, *Theory and Applications of the Analytic Network Process: Decision Making with Benefits, Opportunities, Costs, and Risks*, 3rd ed. Pittsburgh, PA, USA: RWS Publications, 2005.
- [2] T.L. Saaty, "The analytic hierarchy and analytic network measurement processes: Applications to decisions under risk", *European Journal of Pure and Applied Mathematics*, vol. 1, no. 1, pp. 122-196, Jan. 2008. doi: <https://doi.org/10.29020/nybg.ejpam.v1i1.6>.
- [3] Д. Ланде, О. Пучков, та І. Субач, "Система аналізу великих обсягів даних з питань кібербезпеки із соціальних медіа", *Information Technology and Security*, т. 8, № 1, с. 4-18, 2020. doi: <https://doi.org/10.20535/2411-1031.2020.8.1.217993>.
- [4] D. Lande, O. Puchkov, and I. Subach, "Aggregation of information from diverse networks as the basis for training cyber security specialists on processing ultra large data sets", *Information Technology and Security*, vol. 9, no. 1, pp. 4-16, 2021. doi: <https://doi.org/10.20535/2411-1031.2021.9.1.247256>.
- [5] Ю. Кожедуб, С. Василенко, А. Максимець, та В. Гирда, "Концептуальна модель захисту інформації об'єктів критичної інформаційної інфраструктури України", *Information Technology and Security*, т. 9, № 2, с. 151-164, 2021. doi: <https://doi.org/10.20535/2411-1031.2021.9.2.249889>.
- [6] В. Шиповський, "Система показників оцінювання кіберстійкості інформаційних систем об'єктів критичної інфраструктури", *Захист інформації*, т. 25, № 1, с. 37-45, 2023. doi: <https://doi.org/10.18372/2410-7840.25.17597>.
- [7] Д. Юдіна, "Модифікація моделі розрахунку рівня кіберзахисту об'єктів критичної інформаційної інфраструктури", *Кібербезпека: освіта, наука, техніка*, т. 1, № 29, с. 818-836, 2025. doi: <https://doi.org/10.28925/2663-4023.2025.29.943>.
- [8] A. Santos-Olmo et al., "Towards an integrated risk analysis security framework according to a systematic analysis of existing proposals", *Frontiers of Computer Science*, vol. 18, no. 3, art. no. 183808, 2024. doi: <https://doi.org/10.1007/s11704-023-1582-6>.
- [9] W. Xu, and M. Lin, "Information Security Evaluation of Industrial Control Systems Using Probabilistic Linguistic MCDM Method", *Computers, Materials & Continua*, vol. 77, no. 1, pp. 199-222, 2023. doi: <https://doi.org/10.32604/cmc.2023.041475>.
- [10] G. Demir, P. Chatterjee, and D. Pamucar, "Sensitivity analysis in multi-criteria decision making: A state-of-the-art research perspective using bibliometric analysis", *Expert Systems with Applications*, vol. 237, Part C, art. no. 121660, 2024. doi: <https://doi.org/10.1016/j.eswa.2023.121660>.
- [11] B. Efron, and R.J. Tibshirani, *An Introduction to the Bootstrap*. New York, NY, USA: Chapman & Hall, 1994. doi: <https://doi.org/10.1201/9780429246593>.
- [12] W.W. Koczkodaj, "Testing the accuracy enhancement of pairwise comparisons by a Monte Carlo experiment", *Journal of Statistical Planning and Inference*, vol. 69, no. 1, pp. 21-31, 1998. doi: [https://doi.org/10.1016/S0378-3758\(97\)00131-6](https://doi.org/10.1016/S0378-3758(97)00131-6).
- [13] NIST, "Guide to Operational Technology (OT) Security", *NIST Special Publication (SP) 800-82 Rev. 3*, September 2023. doi: <https://doi.org/10.6028/NIST.SP.800-82r3>.

- [14] NIST, “The NIST Cybersecurity Framework (CSF) 2.0”, *NIST Cybersecurity White Paper (CSWP)* 29, Feb. 26, 2024. <https://doi.org/doi: 10.6028/NIST.CSWP.29>.
- [15] Cybersecurity and Infrastructure Security Agency (CISA), “Cross-Sector Cybersecurity Performance Goals”, ver. 1.0.1, March 2023. [Online]. Available: https://www.cisa.gov/sites/default/files/2023-03/CISA_CPG_REPORT_v1.0.1_FINAL.pdf. Accessed on: Oct. 1, 2025.
- [16] The MITRE Corporation, “ICS Matrix”, *MITRE ATT&CK®*, 2025. [Online]. Available: <https://attack.mitre.org/matrices/ics>. Accessed on: Oct. 1, 2025.
- [17] E. Seneta, *Non-negative Matrices and Markov Chains, rev. and expanded 2nd ed.* New York, NY, USA: Springer, 2006. doi: <https://doi.org/10.1007/0-387-32792-4>.
- [18] NIST, “Secure Software Development Framework (SSDF) Version 1.1: Recommendations for Mitigating the Risk of Software Vulnerabilities”, *NIST Special Publication (SP) 800-218*, February 2022. doi: <https://doi.org/10.6028/NIST.SP.800-218>.
- [19] A. Capstick, R. Krishnan, and P. Barnaghi, “AutoElicit: Using Large Language Models for Expert Prior Elicitation in Predictive Modelling”, in *Proc. 42nd Int. Conf. on Machine Learning (ICML 2025)*, Vancouver, BC, Canada, Jul. 13-19, 2025, pp. 6746-6777. [Online]. Available: <https://proceedings.mlr.press/v267/capstick25a.html>. Accessed on: Oct. 1, 2025.
- [20] J.C. González-Quesada, J.R. Trillo, C. Porcel, I.J. Pérez, and F.J. Cabrerizo, “Modelling Large-Scale Group Decision-Making Through Grouping with Large Language Models”, *Future Internet*, vol. 17, no. 9, art. no. 381, 2025. doi: <https://doi.org/10.3390/fi17090381>.
- [21] K.C. Ágoston, and L. Csató, “A lexicographically optimal completion for pairwise comparison matrices with missing entries”, *European Journal of Operational Research*, vol. 314, no. 3, pp. 1078-1086, 2024. doi: <https://doi.org/10.1016/j.ejor.2023.10.035>.
- [22] W.W. Koczkodaj, “A new definition of consistency of pairwise comparisons”, *Mathematical and Computer Modelling*, vol. 18, no. 7, pp. 79-84, 1993. doi: [https://doi.org/10.1016/0895-7177\(93\)90059-8](https://doi.org/10.1016/0895-7177(93)90059-8).
- [23] D.F. Gleich, “PageRank Beyond the Web”, *SIAM Review*, vol. 57, no. 3, pp. 321-363, 2015. doi: <https://doi.org/10.1137/140976649>.
- [24] N. Shbikat, and O.M. Bwaliez, “Enhancing Kendall’s W using genetic algorithm: A computational approach to inter-rater reliability optimization”, *Expert Systems with Applications*, vol. 278, art. no. 127320, Jun. 2025. doi: <https://doi.org/10.1016/j.eswa.2025.127320>.
- [25] B. Xia, T. Bi, Z. Xing, Q. Lu, and L. Zhu, “An Empirical Study on Software Bill of Materials: Where We Stand and the Road Ahead”, in *Proc. 45th Int. Conf. on Software Engineering (ICSE 2023)*, Melbourne, Australia, May 14-20, 2023, pp. 2630-2642. doi: <https://doi.org/10.1109/ICSE48619.2023.00219>.
- [26] A. Krause, and D. Golovin, “Submodular Function Maximization”, in *Tractability: Practical Approaches to Hard Problems*, L. Bordeaux, Y. Hamadi, and P. Kohli, Eds. Cambridge, UK: Cambridge University Press, 2014, pp. 71-104. doi: <https://doi.org/10.1017/CBO9781139177801.004>.
- [27] M. Sviridenko, “A note on maximizing a submodular set function subject to a knapsack constraint”, *Operations Research Letters*, vol. 32, no. 1, pp. 41-43, Jan. 2004. doi: [https://doi.org/10.1016/S0167-6377\(03\)00062-2](https://doi.org/10.1016/S0167-6377(03)00062-2).
- [28] J. Leskovec, A. Krause, C. Guestrin, C. Faloutsos, J. VanBriesen, and N. Glance, “Cost-effective outbreak detection in networks”, in *Proc. 13th ACM SIGKDD Int. Conf. on Knowledge Discovery and Data Mining (KDD 2007)*, San Jose, CA, USA, Aug. 12-15, 2007, pp. 420-429. doi: <https://doi.org/10.1145/1281192.1281239>.
- [29] The MITRE Corporation, “Triton (S1009)”, *MITRE ATT&CK®*. [Online]. Available: <https://attack.mitre.org/software/S1009/>. Accessed on: Oct. 1, 2025.

- [30] Cybersecurity and Infrastructure Security Agency (CISA), “DarkSide Ransomware: Best Practices for Preventing Business Disruption from Ransomware Attacks. Alert AA21-131A”, *Cybersecurity Advisory*, May 2021. [Online]. Available: https://www.cisa.gov/sites/default/files/publications/AA21-131A_Darkside_Ransomware.pdf. Accessed on: Oct. 1, 2025.

REFERENCES

- [1] T.L. Saaty, *Theory and Applications of the Analytic Network Process: Decision Making with Benefits, Opportunities, Costs, and Risks*, 3rd ed. Pittsburgh, PA, USA: RWS Publications, 2005.
- [2] T.L. Saaty, “The analytic hierarchy and analytic network measurement processes: Applications to decisions under risk”, *European Journal of Pure and Applied Mathematics*, vol. 1, no. 1, pp. 122-196, Jan. 2008. doi: <https://doi.org/10.29020/nybg.ejpam.v1i1.6>.
- [3] D. Lande, O. Puchkov, and I. Subach, “System for analysing of big data on cybersecurity issues from social media”, *Information Technology and Security*, vol. 8, iss. 1, pp. 4-18, 2020. doi: <https://doi.org/10.20535/2411-1031.2020.8.1.217993>.
- [4] D. Lande, O. Puchkov, and I. Subach, “Aggregation of information from diverse networks as the basis for training cyber security specialists on processing ultra large data sets”, *Information Technology and Security*, vol. 9, no. 1, pp. 4-16, 2021. doi: <https://doi.org/10.20535/2411-1031.2021.9.1.247256>.
- [5] Y. Kozhedub, S. Vasylenko, A. Maksymets, and V. Hyrda, “Conceptual model of information protection of critical information infrastructure objects of Ukraine”, *Information Technology and Security*, vol. 9, iss. 2, pp. 151-164, 2021. doi: <https://doi.org/10.20535/2411-1031.2021.9.2.249889>.
- [6] V. Shypovskiy, “System of cyber resistance assessment indicators information systems of critical infrastructure objects”, *Information Protection*, vol. 25, iss. 1, c. 37-45, 2023. doi: <https://doi.org/10.18372/2410-7840.25.17597>.
- [7] D. Yudina, “Modification of the model for calculating the level of cyber protection of critical information infrastructure objects”, *Cybersecurity: Education, Science, Technology*, vol. 1, iss. 29, pp. 818-836, 2025. doi: <https://doi.org/10.28925/2663-4023.2025.29.943>.
- [8] A Santos-Olmo et al., “Towards an integrated risk analysis security framework according to a systematic analysis of existing proposals”, *Frontiers of Computer Science*, vol. 18, no. 3, art. no. 183808, 2024. doi: <https://doi.org/10.1007/s11704-023-1582-6>.
- [9] W. Xu, and M. Lin, “Information Security Evaluation of Industrial Control Systems Using Probabilistic Linguistic MCDM Method”, *Computers, Materials & Continua*, vol. 77, no. 1, pp. 199-222, 2023. doi: <https://doi.org/10.32604/cmc.2023.041475>.
- [10] G. Demir, P. Chatterjee, and D. Pamucar, “Sensitivity analysis in multi-criteria decision making: A state-of-the-art research perspective using bibliometric analysis”, *Expert Systems with Applications*, vol. 237, Part C, art. no. 121660, 2024. doi: <https://doi.org/10.1016/j.eswa.2023.121660>.
- [11] B. Efron, and R.J. Tibshirani, *An Introduction to the Bootstrap*. New York, NY, USA: Chapman & Hall, 1994. doi: <https://doi.org/10.1201/9780429246593>.
- [12] W.W. Koczkodaj, “Testing the accuracy enhancement of pairwise comparisons by a Monte Carlo experiment”, *Journal of Statistical Planning and Inference*, vol. 69, no. 1, pp. 21-31, 1998. doi: [https://doi.org/10.1016/S0378-3758\(97\)00131-6](https://doi.org/10.1016/S0378-3758(97)00131-6).
- [13] NIST, “Guide to Operational Technology (OT) Security”, *NIST Special Publication (SP) 800-82 Rev. 3*, September 2023. doi: <https://doi.org/10.6028/NIST.SP.800-82r3>.

- [14] NIST, “The NIST Cybersecurity Framework (CSF) 2.0”, *NIST Cybersecurity White Paper (CSWP)* 29, Feb. 26, 2024. <https://doi.org/doi: 10.6028/NIST.CSWP.29>.
- [15] Cybersecurity and Infrastructure Security Agency (CISA), “Cross-Sector Cybersecurity Performance Goals”, ver. 1.0.1, March 2023. [Online]. Available: https://www.cisa.gov/sites/default/files/2023-03/CISA_CPG_REPORT_v1.0.1_FINAL.pdf. Accessed on: Oct. 1, 2025.
- [16] The MITRE Corporation, “ICS Matrix”, *MITRE ATT&CK®*, 2025. [Online]. Available: <https://attack.mitre.org/matrices/ics>. Accessed on: Oct. 1, 2025.
- [17] E. Seneta, *Non-negative Matrices and Markov Chains, rev. and expanded 2nd ed.* New York, NY, USA: Springer, 2006. doi: <https://doi.org/10.1007/0-387-32792-4>.
- [18] NIST, “Secure Software Development Framework (SSDF) Version 1.1: Recommendations for Mitigating the Risk of Software Vulnerabilities”, *NIST Special Publication (SP) 800-218*, February 2022. doi: <https://doi.org/10.6028/NIST.SP.800-218>.
- [19] A. Capstick, R. Krishnan, and P. Barnaghi, “AutoElicit: Using Large Language Models for Expert Prior Elicitation in Predictive Modelling”, in *Proc. 42nd Int. Conf. on Machine Learning (ICML 2025)*, Vancouver, BC, Canada, Jul. 13-19, 2025, pp. 6746-6777. [Online]. Available: <https://proceedings.mlr.press/v267/capstick25a.html>. Accessed on: Oct. 1, 2025.
- [20] J.C. González-Quesada, J.R. Trillo, C. Porcel, I.J. Pérez, and F.J. Cabrerizo, “Modelling Large-Scale Group Decision-Making Through Grouping with Large Language Models”, *Future Internet*, vol. 17, no. 9, art. no. 381, 2025. doi: <https://doi.org/10.3390/fi17090381>.
- [21] K.C. Ágoston, and L. Csató, “A lexicographically optimal completion for pairwise comparison matrices with missing entries”, *European Journal of Operational Research*, vol. 314, no. 3, pp. 1078-1086, 2024. doi: <https://doi.org/10.1016/j.ejor.2023.10.035>.
- [22] W.W. Koczkodaj, “A new definition of consistency of pairwise comparisons”, *Mathematical and Computer Modelling*, vol. 18, no. 7, pp. 79-84, 1993. doi: [https://doi.org/10.1016/0895-7177\(93\)90059-8](https://doi.org/10.1016/0895-7177(93)90059-8).
- [23] D.F. Gleich, “PageRank Beyond the Web”, *SIAM Review*, vol. 57, no. 3, pp. 321-363, 2015. doi: <https://doi.org/10.1137/140976649>.
- [24] N. Shbikat, and O.M. Bwaliez, “Enhancing Kendall’s W using genetic algorithm: A computational approach to inter-rater reliability optimization”, *Expert Systems with Applications*, vol. 278, art. no. 127320, Jun. 2025. doi: <https://doi.org/10.1016/j.eswa.2025.127320>.
- [25] B. Xia, T. Bi, Z. Xing, Q. Lu, and L. Zhu, “An Empirical Study on Software Bill of Materials: Where We Stand and the Road Ahead”, in *Proc. 45th Int. Conf. on Software Engineering (ICSE 2023)*, Melbourne, Australia, May 14-20, 2023, pp. 2630-2642. doi: <https://doi.org/10.1109/ICSE48619.2023.00219>.
- [26] A. Krause, and D. Golovin, “Submodular Function Maximization”, in *Tractability: Practical Approaches to Hard Problems*, L. Bordeaux, Y. Hamadi, and P. Kohli, Eds. Cambridge, UK: Cambridge University Press, 2014, pp. 71-104. doi: <https://doi.org/10.1017/CBO9781139177801.004>.
- [27] M. Sviridenko, “A note on maximizing a submodular set function subject to a knapsack constraint”, *Operations Research Letters*, vol. 32, no. 1, pp. 41-43, Jan. 2004. doi: [https://doi.org/10.1016/S0167-6377\(03\)00062-2](https://doi.org/10.1016/S0167-6377(03)00062-2).
- [28] J. Leskovec, A. Krause, C. Guestrin, C. Faloutsos, J. VanBriesen, and N. Glance, “Cost-effective outbreak detection in networks”, in *Proc. 13th ACM SIGKDD Int. Conf. on Knowledge Discovery and Data Mining (KDD 2007)*, San Jose, CA, USA, Aug. 12-15, 2007, pp. 420-429. doi: <https://doi.org/10.1145/1281192.1281239>.
- [29] The MITRE Corporation, “Triton (S1009)”, *MITRE ATT&CK®*. [Online]. Available: <https://attack.mitre.org/software/S1009/>. Accessed on: Oct. 1, 2025.

- [30] Cybersecurity and Infrastructure Security Agency (CISA), “DarkSide Ransomware: Best Practices for Preventing Business Disruption from Ransomware Attacks. Alert AA21-131A”, *Cybersecurity Advisory*, May 2021. [Online]. Available: https://www.cisa.gov/sites/default/files/publications/AA21-131A_Darkside_Ransomware.pdf. Accessed on: Oct. 1, 2025.

IGOR SVOBODA
DMYTRO LANDE

PRIORITIZATION OF SECURITY CONTROLS FOR CRITICAL INFRASTRUCTURE USING THE ANALYTIC NETWORK PROCESS AND LARGE LANGUAGE MODELS

This paper addresses the applied problem of prioritizing security controls for national critical infrastructure under state-level threats and severe resource constraints. We use the Analytic Network Process (ANP) to explicitly model nonlinear interdependencies and feedback loops in the system “controls–evaluation criteria–threat vectors–constraints”. A 20-node ANP model is constructed with four clusters: security controls (7 alternatives), evaluation criteria (5), threat vectors (5), and constraints (3). The main novelty is an expert-elicitation workflow based on “virtual experts”. Seven role personas (e.g., ICS engineer, SOC lead, CISO) are instantiated using large language models (LLMs) and used to produce the pairwise judgments required by ANP. The judgments are aggregated with the geometric mean. The resulting inputs demonstrate high consistency (mean Saaty consistency ratio ≈ 0.006 ; mean Koczkodaj index ≈ 0.034), enabling reliable synthesis of the limit supermatrix and global priorities. The final ranking assigns the highest priorities to Network Monitoring and Anomaly Detection (0.1948) and Network Segmentation / Unidirectional Gateways (0.1832), followed by Identity & Privileged Access Management (0.1623), Supply-Chain Security with SBOM and code signing (0.1354), and Incident Response readiness (0.1342). The lowest priority in the considered scenario is Physical Hardening (0.0659). Robustness is confirmed by a leave-one-expert-out (LOEO) analysis and by Monte-Carlo perturbation (1000 trials), which yield stable rankings. Practical usefulness is illustrated with a portfolio selection model under a \$10.2M budget, where a submodular knapsack heuristic selects {Monitoring, Identity, Incident Response, Supply Chain} as the highest-value bundle for threat coverage.

Keywords: analytic network process, cybersecurity, critical infrastructure, state threats, multi-criteria decision-making, large language models, virtual experts.

Свобода Ігор Максимович, аспірант кафедри інформаційної безпеки, Навчально-науковий фізико-технічний інститут Національного технічного університету України “Київський політехнічний інститут імені Ігоря Сікорського”, Київ, Україна, ORCID 0009-0001-1709-2904, i.svoboda@kpi.ua.

Ланде Дмитро Володимирович, доктор технічних наук, професор, завідувач кафедри інформаційної безпеки, Навчально-науковий фізико-технічний інститут Національного технічного університету України “Київський політехнічний інститут імені Ігоря Сікорського”, Київ, Україна, ORCID 0000-0003-3945-1178, dwlande@gmail.com.

Igor M. Svoboda, phd student, department of information security, Institute of physics and technology, National technical university of ukraine “Igor Sikorsky Kyiv Polytechnic Institute”, Kyiv, Ukraine.

Dmytro V. Lande, doctor of technical sciences, professor, head of the department of information security, Institute of physics and technology, National technical university of Ukraine “Igor Sikorsky Kyiv Polytechnic Institute”, Kyiv, Ukraine.

Стаття надійшла до редакції 10.12.2025.

Стаття прийнята до друку після рецензування 01.06.2026.

Дата оприлюднення: 26.06.2026.