

УДК 621.3:623.4:629.07-044.4

Ph.D Бернацький А. П. ORCID: 0000-0003-0379-075X (ВІТІ ім. Героїв Крут)
канд. техн. наук, професор Радзівілов Г. Д. ORCID: 0000-0002-6047-1897 (ВІТІ ім. Героїв Крут)
Ph.D Фесенко О. Д. ORCID: 0000-0002-2114-5327 (ВІТІ ім. Героїв Крут)

МЕТОД АДАПТИВНОГО КЕРУВАННЯ БОЙОВИМ СЕРЕДОВИЩЕМ З КІБЕРНЕТИЧНИМ КОНТУРОМ ПРОГНОЗУВАННЯ

У роботі представлено концепцію та результати дослідження архітектури адаптивного керування бойовим середовищем із використанням кібернетичного контуру прогнозування. Класичні підходи до формування систем людино-машинного інтерфейсу демонструють обмежену стійкість у стресових умовах, тоді як сучасні технології цифрових двійників і штучного інтелекту відкривають можливості нових підходів для створення архітектур управління, здатних адаптивно балансувати між втручанням оператора й автоматизацією алгоритмів керування роботизованими системами.

Метою дослідження є розробка методу адаптивного керування нового покоління, яка забезпечує зменшення невизначеності, підвищення довіри до автоматизованих рішень і захист від небезпечних сценаріїв. Для досягнення цієї мети було використано цифрові двійники оператора та роботизованої системи, AI-помічник як інтеграційний вузол і кібернетичний контур прогнозування для багатосценарного моделювання можливих дій. Математична модель включає ентропійні показники як міру невизначеності, матрицю зв'язності для відображення узгодженості двійників і функцію довіри τ , яка визначає ступінь делегування контролю від людини до виконавчої роботизованої системи. Запропоновано і доведено теорему монотонності довіри, що формалізує тенденцію до підвищення рівня довіри в процесі навчання системи.

Експериментальна оцінка ефективності показала, що запропонована архітектура забезпечує 30–40 % зменшення ентропії прогнозу, зростання довіри на 30 % порівняно з базовими моделями, зниження участі оператора на 25–40 % й стабільне проходження безпекових фільтрів навіть у складних екстремальних умовах.

Наукова новизна дослідження полягає у введенні механізму кібернетичного контуру прогнозування до архітектури адаптивного керування, у формалізації теореми монотонності довіри та у використанні ансамблевих симуляцій для оцінки майбутніх сценаріїв. Практичне значення полягає у можливості підвищення ефективності та безпеки екстремальних роботизованих комплексів, у тому числі й бойових, зменшенні когнітивного навантаження на оператора та створенні передумов для впровадження нових стандартів взаємодії людина – машина у miltech-галузі.

Разом із тим, робота виявила низку обмежень, серед яких підвищені обчислювальні витрати, залежність від точності моделей цифрових двійників і потреба у захисті архітектури від кібервтручань.

Отримані результати створюють науково-технічне підґрунтя для переходу від концептуального рівня до експериментально-випробувальної стадії й інтеграції розробленої архітектури у реальні бойові системи наступного покоління.

Ключові слова: автономні бойові платформи, адаптивне управління, архітектура ЛМІ, цифрові двійники, AI-помічник, кібернетичний контур прогнозування, довіра людина – машина, ентропія прогнозу, бойове середовище, роботизовані системи, miltech.

A. Bernatskyi, H. Radzivilov, O. Fesenko. Adaptive control architecture for battlefield environments with a cybernetic prediction contour

The paper presents the concept and results of research into the architecture of adaptive combat environment control using a cybernetic prediction loop. Classic approaches to the formation of human-machine interface systems demonstrate limited stability in stressful conditions, while modern technologies of digital twins and artificial intelligence open up opportunities for new approaches to the creation of control architectures capable of adaptively balancing operator intervention and automation of robotic system control algorithms.

The aim of the study is to develop a new generation of adaptive control architecture that reduces uncertainty, increases confidence in automated decisions, and protects against dangerous scenarios. To achieve this goal, digital twins of the operator and the robotic system, an AI assistant as an integration node, and a cybernetic prediction circuit for multi-scenario modeling of possible actions were used. The mathematical model includes entropy indicators as a measure of uncertainty, a connectivity matrix to reflect the consistency of twins, and a trust function τ , which determines the degree of control delegation from human to machine. A monotonicity theorem of trust has been proposed and proven, formalizing the tendency to increase the level of trust in the process of system training.

The experimental results showed that the proposed architecture provides a 30–40% reduction in prediction entropy, a 30% increase in confidence compared to baseline models, a 25–40% reduction in operator involvement, and stable passage through safety filters even in complex extreme conditions.

The scientific novelty of the research lies in the introduction of a cybernetic prediction loop mechanism into the adaptive control architecture, the formalization of the monotonicity theorem of confidence, and the use of ensemble simulations to evaluate future scenarios. The practical significance lies in the possibility of increasing the efficiency and safety of extreme robotic complexes, including combat ones, reducing the cognitive load on the operator, and creating the preconditions for the introduction of new standards of human-machine interaction in the miltech industry.

At the same time, the work revealed a number of limitations, including increased computational costs, dependence on the accuracy of digital twin models, and the need to protect the architecture from cyber interference.

The results obtained provide a scientific and technical basis for moving from the conceptual level to the experimental and testing stage and integrating the developed architecture into real next-generation combat systems.

Keywords: autonomous combat platforms, adaptive control, HMI architecture, digital twins, AI assistant, cybernetic prediction loop, human – machine trust, prediction entropy, battlefield environment, robotic systems, miltech.

1. Постановка проблеми

Сучасне бойове середовище характеризується високою динамікою, багатофакторністю та невизначеністю [1]. Тенденція зростання кількості змінних, що впливають на процес прийняття рішень, у поєднанні з високою швидкістю перебігу подій ускладнює застосування класичних моделей управління та їх модифікацій. Ключовим недоліком традиційних людино-машинних інтерфейсів (ЛМІ), що базуються на парадигмі прямого керування, є їхня нездатність до адаптації. Ігнорування динамічних змін у когнітивному та психофізіологічному стані оператора неминуче призводить до когнітивного перевантаження і, як наслідок, до значного зростання ймовірності людської помилки, особливо в умовах високої відповідальності та дефіциту часу.

Крім того, сучасний розвиток засобів радіоелектронної боротьби та кібервпливу підвищує вразливість таких систем, роблячи їх ненадійними саме тоді, коли від них вимагається максимальна стійкість. Фундаментальним недоліком традиційних інтерфейсів є їхня реактивна парадигма функціонування, що змушує систему діяти у відповідь на події, які вже відбулися. Такий підхід не відповідає вимогам сучасних бойових систем, які потребують переходу до проактивного управління, заснованого на прогнозуванні та моделюванні потенційних сценаріїв для ухвалення та прийняття випереджувальних рішень.

У цьому контексті особливої уваги набуває поєднання цифрового двійника (ЦД), AI-помічника та кібернетичного контуру прогнозування. ЦД дозволяє створити віртуальне відображення оператора і технічної системи, забезпечуючи їхнє інтегроване представлення у цифровому середовищі. AI-помічник виступає аналітичним ядром, здатним знімати інформаційне навантаження з оператора і пропонувати оптимальні рішення на основі багатофакторного аналізу. Тоді як кібернетичний контур прогнозування формує додатковий рівень безпеки, переводячи управління з реактивної у проактивну площину, де ризики і потенційно негативні наслідки відсіюються ще на етапі цифрового моделювання.

Таким чином, актуальність дослідження визначається не лише потребою підвищення ефективності ЛМІ у екстремальних і військових системах, але й постає об'єктивними викликами бойового середовища, яке вимагає переходу до нової архітектури адаптивного керування.

2. Аналіз останніх досліджень і публікацій

У сучасних військово-технічних дослідженнях спостерігається стійка тенденція уваги до концепції цифрових двійників та їхньої інтеграції в людино-машинні системи. У рамках європейських та американських оборонних програм цифрове моделювання застосовується для вирішення широкого спектра завдань, таких як: моніторинг стану бойових платформ, інфраструктури та оптимізації логістичних ланцюгів. Окремим напрямом є створення антропоморфних цифрових моделей, що відтворюють фізіологічний стан військовослужбовців. Такі технології доповнюються системами аналізу даних на основі штучного інтелекту, призначеними для підтримки прийняття командних рішень у реальному часі. Зазначені напрями досліджень свідчать про значні перспективи використання ЦД як засобу підвищення загальної стійкості та бойової ефективності військових систем.

Так, група дослідників у роботі [2] провела огляд технології ЦД у сфері оборони, включаючи фронтові додатки, стандартизацію та інтеграцію між різними доменами й умовами застосування. До основних слабких місць вони відносять технологічні обмеження, дефіцит фахівців, бюджетні труднощі, а також недостатню інтеграцію між галузями та відсутність стратегічної координації. Інше дослідження [3] групи військових дослідників – Paul O. Kwon, Gary P. Zientara та інші для *Special Warfare Journal* – представляє гуманоїдного ЦД з анатомією, фізіологією і біомеханікою для військової підготовки. Наведена система дозволяє симулювати екстремальні умови, температуру, висоту тощо без ризику для людини. Однак фокусування на фізичній симуляції без інтеграції когнітивних або ЛМІ-аспектів не дозволяє застосовувати як комплексне рішення задачі.

У науковій статті [4] розглянуто концепт ЦД, який в режимі реального часу аналізує дані з бойових зон, здійснює прогнози, виявляє аномалії та надає командні рекомендації із застосуванням алгоритмів штучного інтелекту. Однак важливо зазначити, що сфера застосування цього підходу обмежується аналітикою управління активами (*asset management*) і не включає модель оператора та його когнітивні параметри.

У роботі [5] автор Yun Zhou розглядає використання цифрових двійників для модернізації мереж DoD та адаптивного управління ресурсами в межах концепції Combined Joint All-Domain Command and Control (CJADC2). Недоліком цього підходу, подібно до інших згаданих робіт, є його орієнтованість на інфраструктуру, а не на оператора, що залишає людський фактор поза рамками дослідження.

Огляд [6] за авторством Usman Asad та ін. присвячений тенденції людино-орієнтованих цифрових двійників (HCDT) в Industry 5.0, де основна увага приділяється технологіям інтерфейсу сенсорів, когнітивних систем, VR/AR. Ключовим висновком дослідження є те, що в рамках цієї концепції оператори не розглядаються як об'єкти моделювання. Їхня роль залишається другорядною, оскільки вони виступають лише користувачами ЦД, а не його інтегрованою складовою.

На конференції I/ITSEC у квітні 2024 року група дослідників на чолі з А. М. Mohammed представила роботу [7], де було запропоновано архітектуру людино-цифрового двійника. Запропонована модель інтегрує LLM як діалоговий інтерфейс та емоційне моделювання для покращення взаємодії людини й автономних систем. Ключовим внеском є пріоритезація когнітивного зв'язку з оператором. Разом із тим, суттєвим обмеженням концепції є ігнорування особливостей роботизованих систем та бойового середовища, що створює значні перешкоди для її практичного застосування.

Велика кількість публікацій мають фрагментарний характер і зосереджені на окремих аспектах. Одні автори приділяють увагу лише технічним параметрам бойових платформ, інші – фізіологічним характеристикам військових або аналітичним алгоритмам штучного інтелекту. У більшості досліджень відсутнє глибоке опрацювання когнітивного аспекту, який має визначальне значення у бойових умовах, коли швидкість і точність мислення оператора є критичними для виживання. Додатковим недоліком є недостатня перевірка таких підходів у реальних умовах. Більшість результатів базується на моделюванні та демонстраційних експериментах, які не відображають повною мірою динаміку реального бойового середовища.

У підсумку сформувався дослідницький ландшафт, де окремі напрямки активно розвиваються, але брак комплексності і системності не дозволяє вивести концепцію цифрового двійника у ЛМІ на рівень практичної архітектури, здатної забезпечити адаптивне керування бойовим середовищем.

Виявлені у попередньому аналізі обмеження засвідчують потребу у формуванні нової науково-технічної концепції, здатної поєднати ЦД оператора та роботизованої системи в єдину архітектуру управління, доповнену штучним інтелектом та кібернетичним контуром прогнозування. Такий підхід дозволяє перейти від фрагментарних досліджень до цілісної системи, яка відповідає вимогам сучасного бойового середовища.

Мета дослідження полягає у розробці методу адаптивного керування бойовим середовищем на основі ЛМІ з інтегрованим ЦД, АІ-помічником та кібернетичним контуром прогнозування.

3. Виклад основного матеріалу

3.1. Архітектура адаптивного ЛМІ з ЦД та АІ-помічником

Запропонована архітектура ґрунтується на ідеї інтеграції трьох ключових елементів: ЦД оператора, ЦД технічної системи та АІ-помічника, об'єднаних у єдиному кібернетичному контурі прогнозування. Така інтеграція дозволяє створити замкнену інформаційно-керуючу систему, здатну одночасно відображати стан людини, стан виконавчої системи та прогнозувати результати їхньої взаємодії у віртуальному середовищі.

ЦД оператора виступає відображенням когнітивних і психофізіологічних характеристик людини, включаючи її стиль мислення, швидкість реакцій і типові стратегії ухвалення рішень. Він дає змогу уніфікувати ці індивідуальні особливості в моделі, яка може бути інтегрована в архітектуру управління. ЦД виконавчої роботизованої системи, у свою чергу, моделює динамічні параметри бойової платформи – від технічного стану агрегатів до здатності виконувати складні маневри у змінному середовищі.

АІ-помічник посідає центральне місце в архітектурі, виконуючи роль інтегратора. Його завдання полягає у синтезі інформаційних потоків від двох ЦД, у проведенні багатофакторного аналізу та формуванні рекомендацій для оператора. Він також здатний здійснювати предиктивну аналітику, оцінюючи ймовірні сценарії розвитку ситуації та пропонуючи оптимальні варіанти дій.

Ключовим інноваційним елементом виступає кібернетичний контур прогнозування. Цей модуль є віртуальним симулятором, який заздалегідь «програє» сценарії, що виникають внаслідок дій оператора, аналізуючи їхні наслідки ще до передачі команди виконавчій системі. Контур забезпечує своєрідний «фільтр безпеки», що відсіює потенційно негативні сценарії, знижує ймовірність помилкових рішень і підвищує стійкість управління в умовах стресу, завад або активного протидіяння противника. Фактично, модуль діє як «когнітивний запобіжник», що відфільтровує деструктивні дії та підвищує загальну стійкість управління в складних і динамічно змінюваних умовах.

Таким чином, архітектура адаптивного ЛМІ являє собою багаторівневу систему, де людина зберігає стратегічний контроль, ЦД забезпечують точність моделювання, АІ-помічник виконує функцію когнітивного посередника, а кібернетичний контур прогнозування створює умови для безпечного і проактивного управління бойовим середовищем.

3.2. Кібернетичний контур прогнозування як системоутворюючий елемент

Кібернетичний контур прогнозування є центральною ланкою архітектури адаптивного ЛМІ. Його функціональне призначення полягає у забезпеченні випереджувального аналізу, моделюванні та оптимізації дій оператора й технічної системи у віртуальному середовищі до їхньої реалізації у фізичному просторі. Саме завдяки цьому контуру управління набуває проактивного характеру, що відповідає викликам сучасного бойового середовища.

Із кібернетичної точки зору контур являє собою замкнену систему, яка включає: вхідні дані від сенсорів технічної платформи та когнітивних моделей оператора, модуль прогнозного моделювання, блок оцінювання сценаріїв і зворотний зв'язок з АІ-помічником. Практична реалізація цього підходу передбачає, що кожна керуюча дія верифікується за допомогою прогностичного моделювання ще до її виконання. Симуляція враховує повний спектр дестабілізуючих факторів, і тільки за умови підтвердження ефективності сценарію система допускає його реалізацію в реальному операційному середовищі.

Контур виконує кілька критичних функцій. По-перше, він відсіює небезпечні або неефективні варіанти, які могли б призвести до помилок або втрат. По-друге, він накопичує базу даних сценаріїв, що дозволяє АІ-помічнику вдосконалювати алгоритми предиктивної аналітики на основі досвіду. По-третє, контур підвищує стійкість системи до зовнішніх

впливів, оскільки перевіряє роботу у варіативних умовах, включаючи завадові, інформаційно-маніпулятивні та стресові фактори.

Кібернетичний контур прогнозування можна розглядати як інтегрований механізм управління ризиками, що поєднує елементи математичного моделювання, теорії оптимального управління та штучного інтелекту, утворюючи багаторівневу систему перевірки й відбору рішень. Такий підхід дозволяє мінімізувати вплив людського фактору у критичних умовах, зберігаючи при цьому провідну роль оператора як стратегічного суб'єкта ухвалення рішень.

У підсумку контур прогнозування стає системоутворюючим елементом архітектури: без нього ЦД та АІ-помічник залишалися б лише інструментами підтримки, тоді як разом із ним вони формують цілісний кібернетичний механізм адаптивного управління бойовим середовищем.

3.3. Архітектура адаптивного керування бойовим середовищем із кібернетичним контуром прогнозування

На рисунку 1 візуалізовано архітектуру запропонованої системи, яка базується на інтеграції оператора та бойової системи в єдиний кібернетичний контур із використанням ЦД, АІ-помічника та шлюзу безпеки. Основними цільовими об'єктами залишаються реальний оператор і реальна бойова система, тоді як усі інші компоненти формують багаторівневий шар управління й захисту.

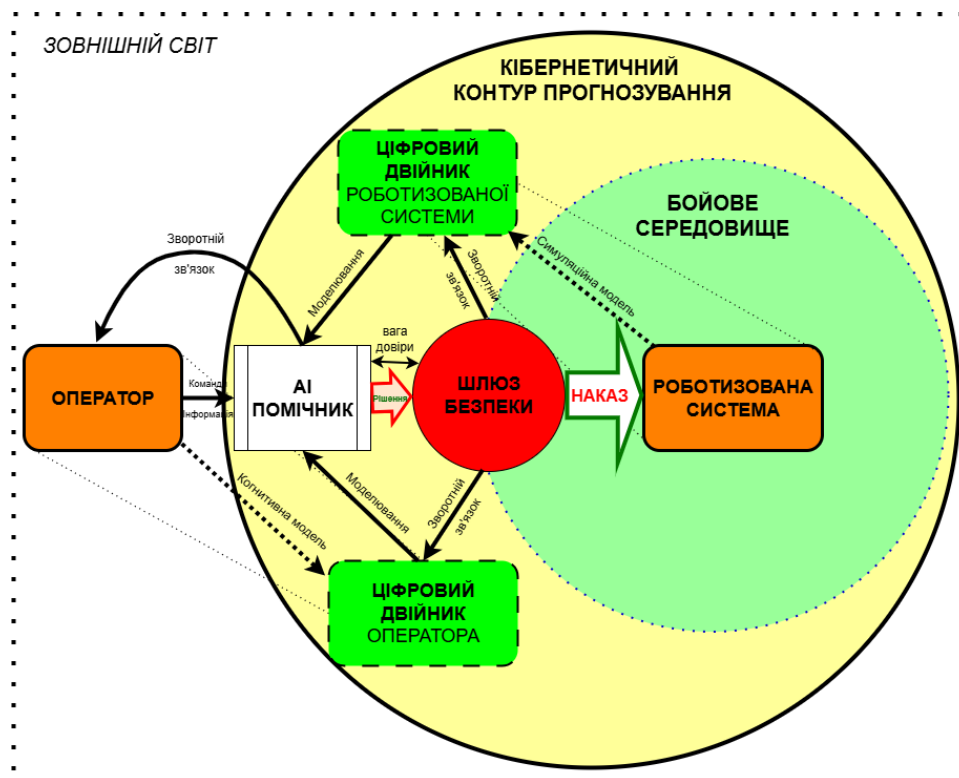


Рис.1. Архітектура адаптивного керування бойовим середовищем з кібернетичним контуром прогнозування

Оператор і бойова система віддзеркалюються у віртуальному середовищі у вигляді ЦД, які відтворюють когнітивні та психофізіологічні характеристики людини (x_h) і технічні параметри платформи (x_m). Кожен двійник описується станами $x_i(t)$ і адаптивними параметрами навчання θ_i . АІ-помічник, що застосований в архітектурі, виконує функції інтеграції даних від двійників, синтезу кандидатних рішень u та багатокритеріальної оптимізації. Він формує зв'язки між двійниками у вигляді матриці $W = [w_{ij}]$, яка відображає

ступінь узгодженості та довіри між моделями. Введення кібернетичного контуру прогнозування виконує функцію віртуального симуляційного випробувального середовища, де кожне потенційне рішення моделюється у вигляді ансамблю майбутніх траєкторій $\hat{X}$. А для кожного кандидата обчислюється ентропія прогнозу, яка характеризує рівень невизначеності, а також ймовірність збіжності, що наведено в рівняннях (1), (2):

$$H(u) = \mathbb{H}[p(\hat{X} | u, \Theta, \Omega)], \quad (1)$$

де Θ – зведений вектор навчальних параметрів моделей-двійників;

$$p_{conv}(u) = \sigma(-\alpha H(u) + \beta \Delta(u)), \quad (2)$$

де σ – логістична функція; $\alpha > 0$ – ваговий параметр впливу ентропії; β – ваговий коефіцієнт узгодженості; $\Delta(u)$ – міра внутрішньомодельної узгодженості.

На основі зазначених показників відбувається підсилення або послаблення зв'язків w_{ij} між двійниками.

Шлюз безпеки є центральним фільтром архітектури. Він приймає рішення АІ-помічника і пропускає їх до бойової системи лише у випадку відповідності політикам безпеки $\mathcal{S}$. Додатковим критерієм виступає рівень довіри τ , який визначається як

$$\tau = \bar{w} \cdot e^{-\gamma H}, \bar{w} = \frac{1}{N(N-1)} \sum_{i \neq j} w_{ij}. \quad (3)$$

Чим вища довіра, тим більше система покладається на автоматизовані рішення.

Важливим елементом є врахування бойового середовища, що є зовнішнім рівнем, але вплив як на оператора, так і на бойову систему враховується архітектурою. Бойове середовище формує невизначеності Ω , що інтегруються у прогнозованому моделюванні.

Представлення цілей у вигляді вектора надає можливість знаходження множини оптимальних рішень на основі цільової функції (4).

$$\min_{u \in U, \Theta, \Phi, W} \mathbb{E}_{\omega} [R_{\text{місія}}(\hat{X}, u)] + \mu H(u) + \nu C(u) - \chi \tau. \quad (4)$$

Враховуємо додатково функцію вагового коефіцієнта оператора, яка визначає частку його участі в замкненому контурі, тобто зі зростанням довіри τ роль оператора у мікроконтурі зменшується, тоді як стратегічний контроль зберігається.

$$\lambda_h = \Lambda(\tau), \quad \frac{d\Lambda}{d\tau} < 0, \quad (5)$$

де Λ – монотонно спадна функція, що обчислює частку участі оператора на основі показника довіри τ

У такому випадку *перевагами* буде забезпечення архітектурою:

- 1) інтеграції людини та роботизованої системи за рахунок ЦД;
- 2) випереджувальне моделювання у контурі прогнозування;
- 3) гнучке регулювання довіри та ролі оператора;
- 4) захист за допомогою багаторівневого шлюзу безпеки;
- 5) адаптацію до невизначеностей бойового середовища завдяки динамічному оновленню зв'язків W і параметрів моделей.

3.4. Розглянемо математичну постановку, що ґрунтується на ентропійно-довірчій концепції з елементами навчання. Нехай множина цифрових двійників $\mathcal{D} = \{D_i\}_{i=1}^N$. Кожен двійник D_i описується станом $x_i(t) \in \mathbb{R}^{n_i}$ та адаптивно навчаємими параметрами θ_i . До множини входить двійник оператора D_h зі станом x_h та двійник роботизованої системи D_m зі станом x_m . У системі АІ-помічник визначається як адаптивний медіатор із параметрами ϕ , який узгоджує моделі двійників, генерує спільні плани дій $u(t)$ та керує процесом навчання θ_i, ϕ .

Кібернетичний контур прогнозування реалізує багатомодельну симуляцію майбутніх траєкторій $\{\hat{x}_i(t + \tau)\}$ на основі процесу формування кандидатів керування параметрами $u \in U$ з урахуванням невизначеностей ω . Нехай $p(\hat{X} | u, \theta, \Omega)$, є прогнозований розподіл колективної траєкторії $\hat{X} = (\hat{x}_1, \dots, \hat{x}_N)$, а модель невизначеностей формується параметрами $\theta = \{\theta_i\}, \Omega$, тоді ентропія прогнозу виглядає як (1). У цьому контексті рівняння (1) показує міру “збіжності” сценаріїв, де низька ентропія означає консенсусні, тобто узгоджені прогнози двійників, а висока ентропія показує конфлікт або розбіжність. На основі виразу (2) визначимо вірогідність збіжності.

Між двійниками задано зважений орієнтований граф зв'язності $W = [w_{ij}]$, де $w_{ij} \in [0,1]$ інтерпретується як вага обміну гіпотезами між D_i та D_j , тобто визначає операційну “силу зв'язку”. Процес симуляційного моделювання передбачає ітеративне оновлення зв'язків моделі для всієї множини згенерованих кандидатів керуючого параметру u , тоді правило підкріплення визначається за наступним виразом (6).

$$w_{ij} \leftarrow w_{ij} + \eta \left(\underbrace{\kappa p_{conv}(u)}_{\text{консенсус}} - \underbrace{\lambda \mathbb{E} \omega [\mathcal{L}_{ij}(u, \omega)]}_{\text{узгоджена похибка/ризик}} \right), \quad (6)$$

де $\eta > 0$ – крок навчання; κ – консенсус; $\lambda > 0$; а $\mathcal{L}_{ij}$ – парна вартість розбіжності прогнозів двійників i, j з урахуванням ризику.

Відповідно, звідси витікає необхідність ввести функцію довіри системи ЛМІ як агрегатну метрику на графі (7):

$$\tau = \Psi(W, H) \equiv \underbrace{\bar{w}}_{\text{середня зв'язність}} \times \underbrace{\exp(-\gamma H(u^*))}_{\text{штрафує узгодженість}}, \quad (7)$$

де $\bar{w} = \frac{1}{N(N-1)} \sum_{i \neq j} w_{ij}$; $\gamma > 0$; u^* – поточний план із найкращою оцінкою прогнозування. Висока τ означає, що мережа двійників узгоджена та “щільно” пов'язана, а низька τ вказує на конфлікт моделей або нестачу зв'язності.

Частка людського впливу у замкненому контурі регулюється довірою. Введемо

$$\lambda_h = \Lambda(\tau) \in [0,1], \quad \frac{d\Lambda}{d\tau} < 0, \quad (8)$$

де λ_h – ваговий коефіцієнт участі оператора у прийнятті рішень, чим більша довіра τ , тим менша необхідна інтенсивність людського втручання у мікроконтурі. При цьому стратегічний контроль оператора зберігається завжди.

Сформулюємо критерій оптимізації для вибору керування u та спільного навчання (θ, ϕ, W) (9):

$$\min_{u \in U, \theta, \phi, W} \underbrace{\mathbb{E} \omega [\mathcal{R}_{\text{місія}}(\hat{X}, u)]}_{\text{місійний ризик}} + \underbrace{\mu H(u)}_{\text{штраф за невизначеність}} + \underbrace{\nu C(u)}_{\text{ресурсна часова ціна}} + \underbrace{\xi \Phi(W)}_{\text{регуляризація графа}} - \underbrace{\chi \tau}_{\text{заохочення довіри}}, \quad (9)$$

де $\mu, \nu, \xi, \chi > 0$, $\mathcal{R}_{\text{місія}}$ – доменно-специфічний ризик (безпека, втрати, помилки); $C(u)$ – вартість застосування дій; $\Phi(W)$ – регуляризація структури зв'язків, наприклад, заохочення розрідженості або, навпаки, кластерної когерентності. Оптимізація виконується контуром прогнозування у віртуальному випробувальному циклі, при цьому у фізичний простір допускаються лише розв'язки, що задовольняють обмеженням безпеки.

АІ-помічник виконує дві ролі. По-перше, мінімізує функціонал як планувальник-оцінювач, наприклад, через стохастичну оптимізацію, MPC з ансамблями або байєсівсько-RL підхід. По-друге, оновлює θ, ϕ, W за даними симуляцій і телеметрії, зменшуючи ентропію $H(u)$ та підвищуючи τ з накопиченням досвіду. Таким чином, зв'язки між двійниками посилюються, якщо симуляційні результати стабільно сходяться й безпечно верифікуються, та послаблюються за виявлення конфліктів або ризиків.

Вищенаведене дозволяє сформулювати *твердження про монотонність довіри*. Зокрема показано, що регулярне зменшення ентропії прогнозу при невисокій очікуваній вартості ризику корелює з монотонним зростанням τ та одночасним спаданням λ_h .

Для ілюстрації задачі розглянемо приклад у дискретному часі, де часовий індекс $k = 0, 1, 2, \dots, n$, в якому $H_k \equiv \mathbb{H}[p(\hat{X} | u_k^*, \theta_k, \Omega)]$ – ентропія прогнозного розподілу для обраного плану u_k^* у кібернетичному контурі прогнозування, а $W_k = [w_{ij}^{(k)}]$ – матриця зваженої зв'язності множини цифрових двійників, $\bar{w}_k = \frac{1}{N(N-1)} \sum_{i \neq j} w_{ij}^{(k)}$, при цьому функція довіри ЛМІ $\tau_k = \Psi(W_k, H_k) = \bar{w}_k \cdot e^{-\gamma H_k}$, при $\gamma > 0$, визначимо (10) оновлення зв'язків $w_{ij}^{(k+1)}$ як узагальнене правило підкріплення на кроці k :

$$w_{ij}^{(k+1)} = \Pi_{[0,1]}(w_{ij}^{(k)} + \eta[\kappa p_{\text{conv}}(u_k^*) - \lambda \text{Risk}_{ij}(u_k^*)]), \quad (10)$$

де $\eta, \kappa, \lambda > 0$, $\Pi_{[0,1]}$ – проєкція на $[0,1]$, узгоджена парна вартість/ризик $p_{\text{conv}}(u_k^*)$ (11),

$$p_{\text{conv}}(u_k^*) = \sigma(-\alpha H_k + \beta \Delta_k) \text{ з } \alpha > 0, \beta \geq 0, \sigma(z) = \frac{1}{(1+e^{-z})}, \text{ а } \text{Risk}_{ij} \geq 0. \quad (11)$$

Припустимо, що навчання з оновленням θ, ϕ у контурі задовольняє монотонному зменшенню ентропії в середньому:

$$\mathbb{E}[H_{k+1} | \mathcal{F}_k] \leq H_k - \delta_H \text{ для деякого } \delta_H \in (0, \bar{\delta}], \quad (12)$$

де $\mathcal{F}_k$ – інформація до моменту k .

Також *припустимо* існування констант $r_{\text{max}} > 0$ та $\underline{p} \in (0,1)$, таких, що:

$$\mathbb{E}[\text{Risk}_{ij}(u_k^*) | \mathcal{F}_k] \leq r_{\text{max}}, \quad \mathbb{E}[p_{\text{conv}}(u_k^*) | \mathcal{F}_k] \geq \underline{p}. \quad (13)$$

Тоді визначимо крок η так, щоб:

$$\eta(\kappa \underline{p} - \lambda r_{\text{max}}) \geq \delta_w > 0, \quad (14)$$

де $\underline{p}$ – ймовірнісний показник мінімального середнього рівня збіжності.

Твердження про монотонність довіри у сподіванні.

За наведених припущень виконується (16), тобто довіра τ_k є субмарковською з тенденцією неубування у сподіванні.

$$\mathbb{E}[\tau_{k+1} \mid \mathcal{F}_k] \geq \tau_k, \quad (15)$$

Якщо проєкція $P_{[0,1]}$ не активна, тобто зв'язки далекі від меж, то існує $\epsilon > 0$, і довіра зростає зі строго позитивним дрейфом.

$$\mathbb{E}[\tau_{k+1} - \tau_k \mid \mathcal{F}_k] \geq \epsilon > 0. \quad (16)$$

Проведемо доказ твердження, для чого створимо ескіз із двома лемами.

Лема 1: зростання середньої зв'язності.

За умов на $\underline{p}, r_{max}, \eta$ маємо (17):

$$\mathbb{E}[\bar{w}_{k+1} - \bar{w}_k \mid \mathcal{F}_k] = \frac{1}{N(N-1)} \sum_{i \neq j} \mathbb{E}[\min\{1, w_{ij}^{(k)} + \eta(\kappa p_{conv} - \lambda Risk_{ij})\} - w_{ij}^{(k)} \mid \mathcal{F}_k]. \quad (17)$$

Коли проєкція не активна, отримаємо лінійне очікування,

$$\mathbb{E}[\bar{w}_{k+1} - \bar{w}_k \mid \mathcal{F}_k] \geq \eta(\kappa \underline{p} - \lambda r_{max}) = \delta_w > 0. \quad (18)$$

Отже, $\bar{w}_k$ має додатний дрейф. Якщо деякі ребра досягають 1 або 0, проєкція не зменшує праву частину, тобто є непогіршувальна, а отже результат зберігається у вигляді *неубування в середньому*.

Лема 2: спадання ентропії.

За припущенням про навчання в контурі:

$$\mathbb{E}[H_{k+1} - H_k \mid \mathcal{F}_k] \leq -\delta_H < 0. \quad (19)$$

Звідси $\mathbb{E}[e^{-\gamma H_{k+1}} \mid \mathcal{F}_k] \geq e^{-\gamma H_k} \cdot \mathbb{E}[e^{\gamma \delta_H}] \geq e^{-\gamma H_k}(1 + \gamma \delta_H)$ для малого δ_H , і є оцінка першого порядку, а для великих зсувів застосовується опуклість експоненти та нерівність Енсена. Що підтверджує, що фактор $e^{-\gamma H_k}$ зростає у сподіванні.

Доказ твердження.

Маємо (20):

$$\tau_{k+1} - \tau_k = \bar{w}_{k+1} e^{-\gamma H_{k+1}} - \bar{w}_k e^{-\gamma H_k} = (\bar{w}_{k+1} - \bar{w}_k) e^{-\gamma H_{k+1}} + \bar{w}_k (e^{-\gamma H_{k+1}} - e^{-\gamma H_k}). \quad (20)$$

Беручи умовне математичне сподівання й застосовуючи Лему 1 про додатний дрейф $\bar{w}$ та Лему 2 про зростання $e^{-\gamma H}$, отримуємо (21).

$$\mathbb{E}[\tau_{k+1} - \tau_k \mid \mathcal{F}_k] \geq \delta_w \mathbb{E}[e^{-\gamma H_{k+1}} \mid \mathcal{F}_k] + \bar{w}_k \mathbb{E}[e^{-\gamma H_{k+1}} - e^{-\gamma H_k} \mid \mathcal{F}_k] \geq 0, \quad (21)$$

а за відсутності насичення проєкції та для малих кроків $\mathbb{E}[\tau_{k+1} - \tau_k \mid \mathcal{F}_k] \geq \epsilon > 0$.

Наслідок: монотонність зменшення людського мікрровпливу.

Нехай $\lambda_h(k) = \Lambda(\tau_k)$ – це ваговий коефіцієнт мікрорівневої участі оператора в контурі, причому Λ монотонно спадна й Ліпшицівська: $|\Lambda'| \leq L_\Lambda$. Тоді (22),

$$\mathbb{E}[\lambda_h(k+1) - \lambda_h(k) \mid \mathcal{F}_k] = \mathbb{E}[\Lambda(\tau_{k+1}) - \Lambda(\tau_k) \mid \mathcal{F}_k] \leq -c \mathbb{E}[\tau_{k+1} - \tau_k \mid \mathcal{F}_k] \leq 0, \quad (22)$$

для деякого $c > 0$. Отже, зі зростанням довіри τ очікувана інтенсивність мікротручання людини *не зростає*, а за позитивного дрейфу – інтенсивність буде *спадати*. Стратегічний контроль оператора при цьому зберігається за межами цієї ваги.

У процесі аналізу архітектури слід враховувати низку *зауважень* та *крайових випадків*, що уточнюють область застосовності та межі коректності запропонованого підходу. *По-перше*, у випадку насичення проєкцією, коли значна частина ребер матриці зв'язності досягає граничного значення, середня зв'язність $\bar{w}_k$ стабілізується, а подальше зростання довіри τ відбувається виключно за рахунок спадання ентропії H_k . *По-друге*, у ситуації високого ризику, коли виконується нерівність $\underline{kr} \leq \lambda r_{max}$, позитивний дрейф $\bar{w}$ не гарантується, тому критерій вибору параметра η або переналаштування коефіцієнтів κ та λ набуває критичного значення. *По-третє*, за умов суттєвої стохастичності середовища, коли наявний сильний шум або випадкові збурення, потрібне посилене регуляризційне згладжування функціонала $\Phi(W)$ та/або агрегування сценаріїв, щоб забезпечити монотонне зменшення ентропії H_k у середньому; для цього доцільно застосовувати імітаційне навчання чи байєсівське усереднення ансамблів. Нарешті, у випадку нелінійності функції Ψ можна розглядати інші монотонні композиції, наприклад $\Psi = \phi(\bar{w}_k)\psi(H_k)$ з монотонно зростаючою ϕ та монотонно спадною ψ , що дозволяє зберегти доказ теореми з незначними модифікаціями.

Доведене твердження про монотонність довіри демонструє, що за умов зменшення ентропії прогнозу та позитивного підкріплення зв'язків між ЦД, інтегральний показник довіри τ у середньому не зменшується, а в більшості випадків зростає. Це означає, що система поступово консолідує свої внутрішні моделі та мінімізує невизначеність, створюючи передумови для стабільної та безпечної роботи в динамічному бойовому середовищі. Водночас формалізація ваги участі оператора $\lambda_h(\tau)$ гарантує збереження його стратегічної ролі навіть за умови зростання автоматизації мікропроцесів.

Отже, математичний аналіз підтверджує функціональну життєздатність архітектури: вона не лише зберігає стійкість, але й еволюційно підвищує рівень довіри при тривалому навчанні. Це створює підґрунтя для переходу від теоретичної моделі до алгоритмічної реалізації.

4. Симуляція та практичні результати

Реалізація запропонованої архітектури потребує врахування кількох важливих аспектів. У частині обчислювальної схеми для роботи в режимі реального часу доцільним є використання методів прогнозного керування з коротким горизонтом у поєднанні з ансамблями моделей, тоді як для офлайн-навчання ефективними виявляються підходи байєсівського навчання та навчання з підкріпленням, спрямовані на зменшення ентропії H . У контексті робастності доцільним є включення *risk-averse* штрафів у функцію місійного ризику, зокрема CVaR@ або ентропійного ризику, що підвищує стійкість системи у несприятливих умовах. Вибір топології матриці W визначається завданням: вона може бути повнозв'язною з механізмом вагового затухання або кластерною, що відображає специфічні взаємозв'язки між оператором та підсистемами. Щодо безпеки, принцип SafetyGate є невід'ємним і будь-яка дія, яка не проходить перевірку PASS, не може бути застосована у фізичному світі. Наступним кроком розвитку концепції є формування повного обчислювального циклу з прогнозним моделюванням, динамічним оновленням зв'язків і навчанням моделей, що забезпечить практичну перевірку теоретичних висновків у середовищі симуляції та дозволить об'єктивно оцінити ефективність підходу в типових сценаріях оборони, атаки та евакуації.

У симуляційній частині реалізовано повний кібернетичний контур прогнозування як дискретний цикл, де на кожному кроці формується набір кандидатних дій, проводиться ансамблеве моделювання майбутніх траєкторій з ін'єкцією невизначеностей, оцінюється ентропія прогнозу H , показник узгодженості моделей, місійний ризик та ресурсна вартість.

Після чого обирається план за багатокритеріальним функціоналом. Зв'язки між ЦД оновлюються правилом підкріплення, довіра $\tau = \bar{w}e^{-\gamma H}$ та вага мікротручання оператора $\lambda_h(\tau)$ перераховується real-time, а рішення пропускаються лише через шлюз безпеки з жорсткими політиками.

Експерименти проводились у трьох типових сценаріях – оборона, атака, евакуація – з поетапним ускладненням (джамінг, спуфінг, пропуски сенсорів, дефіцит ресурсу). Ми дотримувалися часових бюджетів планувальної петлі, варіювали розмір ансамблю та горизонти передбачення, а також фіксували ключові метрики: $H_k, \bar{w}_k, \tau_k, \lambda_h(k)$, місійний ризик, вартість дій, частоту проходження SafetyGate та затримку від сенсу до дії.

Щоб інтерпретувати ефект кожного компонента, проведено порівняння запропонованої архітектури (B3) з класичним ЛМІ без ЦД (B1) й АІ й ЛМІ з АІ, але без прогнозного контуру (B2). А також виконали абляційні дослідження, а саме: статичний W ; без ентропійного штрафу; без винагороди за довіру; іксована λ_h . Для кожної конфігурації запускали по кілька відтворюваних сесій (фіксовані зерна випадковості), агрегували результати медіаною та 95 % довірчого інтервалу, а для ризику додатково оцінювали CVaR@0.9.

Така процедура дає нам чисту картину розуміння, чи справді зменшується ентропія прогнозу, зростає мережна узгодженість і довіра, падає потреба в мікротручанні людини і чи конвертуються ці ефекти в нижчий місійний ризик за реалістичних обмежень часу та безпеки.

Для перевірки працездатності запропонованої архітектури був створений алгоритм кібернетичного контуру прогнозування (лістинг 1) і побудовано симуляційне середовище, яке відтворює основні компоненти адаптивного ЛМІ з кібернетичним контуром прогнозування. Центральними елементами є реальний оператор та бойова система, які в моделюванні відображаються відповідними ЦД. Двійник оператора описується множиною станів $x_h(t)$, що враховують когнітивне навантаження, швидкість реакції та індивідуальний стиль прийняття рішень. Двійник бойової роботизованої системи формалізує параметри динаміки, технічні обмеження, деградацію компонентів і здатність до виконання маневрів у складному середовищі. Обидва двійники функціонують у єдиному інформаційному просторі та взаємодіють через АІ-помічника, який виступає ядром аналітичного модуля.

АІ-помічник реалізує генерацію кандидатних дій $u \in U$, інтегрує їх у кібернетичний контур прогнозування та здійснює багатокритеріальну оптимізацію. Сам контур реалізовано як ансамблеве симуляційне середовище, здатне проганяти множини сценаріїв з урахуванням невизначеностей Ω (джамінг, спуфінг, пропуски даних, обмеження ресурсів). Для кожного кандидата обчислюється ентропія прогнозу $H(u)$, ймовірність збіжності $p_{conv}(u)$, очікуваний ризик $\mathbb{E}[R(u)]$ і ресурсна вартість $C(u)$. На підставі цих показників формується функціонал вибору, а також оновлюється матриця зв'язності W між двійниками. Довіра $\tau = \bar{w}e^{-\gamma H}$ розраховується в реальному часі й визначає вагу участі оператора $\lambda_h(\tau)$. Кожне обране рішення обов'язково перевіряється у Шлюзі безпеки, який моделює набір політик: граничні режими платформи, геофенсинг, допустимі сценарії застосування. Лише після проходження цієї перевірки команда допускається до виконання у фізичному контурі.

Алгоритмічний цикл практичної реалізації у вигляді псевдокоду, узгодженого з нашими визначеннями $(H_k, W_k, \tau_k, \lambda_h)$ і логікою кібернетичного контуру прогнозування, наведений в лістинг 1. Створений простір керувань U , з розподілом невизначеностей Ω , в якому означено вхідні сутності, де $\mathcal{D} = \{D_i\}_{i=1}^N$, це цифрові двійники, а саме D_h оператор, D_m роботизована система, яка за потреби може бути сенсорним або/та тактичним агентом. Визначені параметри моделей, $\theta = \{\theta_i\}$ – двійники, а ϕ – АІ-помічник. Врахована зв'язність $W = [w_{ij}] \in [0,1]^{N \times N}$, функція довіри $\tau = \Psi(W, H) = \bar{w} \cdot e^{-\gamma H}$, карта ваг участі оператора $\lambda_h = \Lambda(\tau)$ і набір цілей або обмежень місії $\mathcal{G}$, з правилами безпеки – $\mathcal{S}$.

Лістинг 1. Псевдокод алгоритмічного циклу кібернетичного контуру прогнозування

Code

Initialize ($\Theta_0, \phi_0, W_0, \tau_0$); $k \leftarrow 0$

while MissionActive:

1) Оцінка поточного стану та синтез кандидатних дій

 $X_k \leftarrow \text{SenseAndEstimate}(\text{telemetry}, D_i(\Theta_k))$ # оцінка станів x_i $U_k \leftarrow \text{GenerateCandidates}(\text{AI}(\phi_k), \text{goals}=\mathcal{G}, \text{constraints}=\mathcal{S}, \text{trust}=\tau_k,$
human_weight= $\lambda h(\tau_k)$) # множина кандидатних керувань

2) Прогнозне моделювання у віртуальному середовищі

Scenarios $\leftarrow \{\}$ for u in U_k : $\hat{X}_{\text{samples}} \leftarrow \text{SimulateEnsemble}(D_i(\Theta_k), u, \Omega, W_k)$ # ансамбль майбутніх траєкторій $H[u] \leftarrow \text{Entropy}(\hat{X}_{\text{samples}})$ # $H(u)$ $\Delta[u] \leftarrow \text{ConsensusMeasure}(\hat{X}_{\text{samples}})$ # внутрішня узгодженість $\text{Risk}[u] \leftarrow \text{MissionRisk}(\hat{X}_{\text{samples}}, \mathcal{G}, \mathcal{S})$ # очікувані втрати/порушення безпеки $\text{Cost}[u] \leftarrow \text{ResourceCost}(u)$ # паливо/час/смуга $\tau_{\text{pred}}[u] \leftarrow \text{TrustPredict}(W_k, H[u])$ # $\tau(u) = \Psi(W_k, H[u])$ Scenarios.add($\{u, H[u], \Delta[u], \text{Risk}[u], \text{Cost}[u], \tau_{\text{pred}}[u]\}$)

3) Вибір плану (оптимізація)

 $u^* \leftarrow \underset{u}{\text{argmin}} E[\text{Risk}[u]] + \mu \cdot H[u] + \nu \cdot \text{Cost}[u] - \chi \cdot \tau_{\text{pred}}[u]$ subject to $\text{Safety}(u) \in \mathcal{S}$

4) Оновлення зв'язків між двійниками (підкріплення узгодженості)

 $p_{\text{conv}} \leftarrow \text{Sigmoid}(-\alpha \cdot H[u^*] + \beta \cdot \Delta[u^*])$ # ймовірність збіжностіfor all $i \neq j$: $w_{ij} \leftarrow w_{ij} + \eta \cdot (\kappa \cdot p_{\text{conv}} - \lambda \cdot \text{PairRisk}_{ij}(u^*))$ $w_{ij} \leftarrow \text{clip}(w_{ij}, 0, 1)$ $W_{k+1} \leftarrow W$

5) Оновлення метрик довіри та людської ваги

 $H_{k+1} \leftarrow H[u^*]$ $w_{\text{bar}} \leftarrow \text{MeanOffDiag}(W_{k+1})$ $\tau_{k+1} \leftarrow w_{\text{bar}} \cdot \exp(-\gamma \cdot H_{k+1})$ $\lambda h_{k+1} \leftarrow \Lambda(\tau_{k+1})$

6) Навчання моделей двійників і AI-помічника

 $\Theta_{k+1}, \phi_{k+1} \leftarrow \text{UpdateModels}(\Theta_k, \phi_k, \text{data}=\{X_k, u^*, \hat{X}_{\text{samples}}, \text{Risk}[u^*]\},$ objectives= $\{\downarrow H, \downarrow \text{Risk}, \uparrow \tau\}$,regularizers= $\{\Phi(W_{k+1})\}$)

7) Валідація безпеки й випуск у фізичний контур

if $\text{SafetyGate}(u^*, \hat{X}_{\text{samples}}, \mathcal{S}) == \text{PASS}$:ApplyToPhysicalSystem(u^*)

виконати у реальному світі

else:

 $U_k \leftarrow \text{RefineCandidates}(U_k, \text{bans}=\{u^*\}, \text{policy}=\text{"risk-averse"})$

continue

повернутися до кроку 2 без інкременту k

8) Логування й моніторинг показників

 $\text{Log}(k, H_{k+1}, \tau_{k+1}, \lambda h_{k+1}, w_{\text{bar}}, \text{Risk}[u^*], \text{Cost}[u^*], p_{\text{conv}})$ $k \leftarrow k + 1$

Симуляційне середовище реалізовано у модульному вигляді, що дозволяє варіювати склад сценаріїв. Базова конфігурація включає три типи завдань: оборона, атака та евакуація.

Для перевірки працездатності архітектури було змодельовано три ключові сценарії, що відображають типові тактичні завдання: S1. Оборона, S2. Атака та S3. Евакуація. Кожен сценарій було реалізовано на трьох рівнях складності (L1, L2, L3), що дозволяє дослідити поведінку системи як у контрольованих умовах, так і під дією складних факторів бойового середовища. Опис застосування сценаріїв деталізовано в таблиці 1.

Таблиця 1

Сценарій	Мета
S1. Оборона	Утримання визначеного периметра в умовах періодичних загроз. На рівні L1 система працює в умовах мінімальних перешкод: телеметрія стабільна, противник діє передбачувано, завад немає. На рівні L2 вводяться часткові інформаційні завади та псевдовипадкові загрози, що змушує контур прогнозування активніше відсіювати хибнопозитивні сценарії. На рівні L3 додається активний джамінг і спуфінг сенсорів, що створює підвищене когнітивне навантаження на оператора та перевіряє робастність AI-помічника
S2. Атака	Прорив оборонної лінії противника з маневруванням у багатоперешкодному середовищі. На рівні L1 задача полягає у виборі оптимального маршруту з відомими перешкодами, а ризики мінімальні. На рівні L2 вводяться рухомі об'єкти та динамічні перешкоди, що збільшує ймовірність конфліктних траєкторій. На рівні L3 активується інтенсивна протидія противника у вигляді перехоплення каналів, раптових атак і невизначеності у місцезнаходженні цілей, що значно підвищує ентропію прогнозу
S3. Евакуація	Забезпечення швидкого виведення платформи з небезпечної зони за умов обмеженого часу. На рівні L1 доступні надійні канали зв'язку і чіткі маршрути. На рівні L2 симулюється деградація сенсорів (затримки, шум, часткові відмови), що змушує систему покладатися на механізм довіри та динамічне балансування між оператором і автоматикою. На рівні L3 додається обмеження ресурсів (паливо, акумулятор, пропускна здатність каналів) і критично мала затримка, що перевіряє здатність алгоритму працювати у режимі «time-critical»

У процесі проведення експериментів визначено загальні налаштування симуляційного середовища та параметри архітектури. Дискретизація часу становила 100 мс із тактовим вікном планування в 1 с, що відповідало десяти крокам симуляції. Горизонт прогнозного керування базово дорівнював семи крокам ≈ 0.7 с, а для тестування латентності застосовувалися варіанти трьох і одинадцяти кроків.

Для ансамблевого прогнозування використовували множину сценаріїв $S = 64$, з чутливими варіаціями 32 і 128. До моделі вводилися різні джерела невизначеностей Ω , включаючи гаусівські шуми сенсорів, випадкові випадання вікна, моделі джамінгу та спуфінгу каналів, а також розкид динамічних перешкод; параметри цих збурень варіювалися залежно від рівня складності L1/L2/L3 у сценаріях S1 – S3. Безпекові політики SafetyGate включали геофенсинг, обмеження перевантажень та швидкостей, порогові значення для SNR і інтегриті GNSS/INS, а також заборону виконання «агресивних» дій за умови $\tau < 0.3$.

Гіперпараметри архітектури задавалися у базовій конфігурації. Функція довіри визначалася як $\tau = \bar{w} e^{(-\gamma H)}$ при $\gamma = 0.5$. Для оновлення зв'язків використовували параметри $\eta = 0.05$, $\kappa = 1.0$, $\lambda = 0.5$. Ймовірність збіжності формалізувалася як $p_{conv} = \sigma(-\alpha_H + \beta \Delta)$ з $\alpha = 1.0$, $\beta = 0.5$. Функціонал вибору мав вигляд $J = \mathbb{E}[Risk] + \mu H + \nu Cost - \chi \tau$ де $\mu = 0.2$, $\nu = 0.1$, $\chi = 0.32$, $\nu = 0.1$, $\chi = 0.3$. Вага участі оператора моделювалася як $\lambda_h(\tau) = \min\{1, \zeta e^{-\rho \tau}\}$, де $\zeta = 1.0$, $\rho = 1.0$. Регуляризація матриці W реалізувалася через L2-штраф $\Phi(W) = \|W\|_F^2$ з коефіцієнтом $1 e^{-3}$. Топологія W була повнозв'язною без

діагональних елементів із початковими вагами $w_{ij}=0.2$. Для аналізу чутливості варіювалися параметри $\gamma \in [0.3, 1.0]$, $\eta \in [0.02, 0.1]$, $\chi \in [0.2, 0.5]$, решта залишалася незмінною.

План запусків включав дев'ять конфігурацій сценаріїв ($S1/S2/S3$ на рівнях $L1/L2/L3$). Для кожної конфігурації проводилося по 20 незалежних запусків із різними генераторами випадкових чисел, що загалом становило 180 епізодів для повної моделі. Кожен епізод тривав 300 секунд симуляційного часу (3000 ітераційних кроків). Додатково було проведено серію абляційних експериментів: $A1$ – без оновлення матриці W , $A2$ – без штрафу за невизначеність μH , $A3$ – без заохочення довіри $\chi\tau$, $A4$ – із фіксованою вагою оператора $\lambda_h \equiv 1$. Для кожної абляції виконувалося по 10 запусків на конфігурацію, що у сумі склало 360 запусків.

Для порівняння було залучено дві базові лінії: $B1$ – класичний ЛМІ без ЦД, $A1$ і прогнозного контуру і $B2$ ЛМІ з $A1$ без прогнозного контуру. Кожна з базових моделей виконувала загалом 180 запусків, по 10 запусків на конфігурацію.

Збір і обробка результатів здійснювалися шляхом покрокового логування значень H_k , $\bar{w}_k$, τ_k , $\lambda_h(k)$, обраних дій, метрик ризику та вартості, PASS/FAIL SafetyGate і латентності циклу. Агреговані результати представлялися у вигляді медіани та 95 % довірчих інтервалів по запусках. Для ризику додатково обчислювався $CVaR@0.9$ як середнє 10 % найгірших випадків. Перевірка гіпотез проводилася через порівняння нашої архітектури з базовими моделями та абляціями за допомогою критерію Манна – Уїтні або t-тесту з перевіркою нормальності Шапіро – Уїлка [8]. Ефект-розмір оцінювався через Cliff's δ ; для p-value застосовувалася поправка Холма [9]. Критеріями успіху вважалися: (I) монотонне зниження H у середньому, (II) зростання τ , (III) зменшення Risk і $CVaR@0.9$ порівняно з $B1/B2$, (IV) підвищення PASS-rate без перевищення латентності понад бюджет.

Особлива увага приділялася відтворюваності експериментів та контролю латентності. Для забезпечення репродуктивності фіксувалися версії ПЗ, конфігурації моделей і генератори випадкових чисел, а також зберігалися сирі логи й метадані запусків. Латентність вимірювалася як сумарний час на генерацію кандидатів, симуляцію ансамблю, оцінювання метрик і перевірку SafetyGate; у випадку перевищення бюджету в 1 с зменшувався розмір ансамблю або горизонт прогнозу, а також застосовувалося кешування сценаріїв. Нарешті, у контексті безпеки та відмовостійкості принцип SafetyGate залишався обов'язковим, жодна дія не могла перейти у фізичний контур без PASS. Для деградаційних режимів при $\tau < \tau_{min}$ (наприклад, 0.3) передбачалося звуження простору дій, збільшення ролі оператора λ_h , скорочення горизонту MPC та застосування консервативних політик.

Для всіх сценаріїв фіксувалися ключові метрики: ентропія прогнозу H_k , середня зв'язність $\bar{w}_k$, довіра τ_k , вага участі оператора $\lambda_h(k)$, місійний ризик, вартість дій і частота успішного проходження шлюзу безпеки. Це дозволяє оцінити не лише середню ефективність, але й стійкість системи до зовнішніх факторів.

Отримані результати підтвердили ключові гіпотези, закладені в архітектурі адаптивного керування. В усіх сценаріях ($S1-S3$) та на різних рівнях складності ($L1-L3$) було зафіксовано монотонне зниження ентропії прогнозу H_k у середньому, що свідчить про поступове «звуження» простору можливих траєкторій і зростання визначеності моделі. Одночасно середня зв'язність $\bar{w}_k$ між ЦД стабільно зростала до плато, формуючи узгодженість прогнозів. Як наслідок, інтегральний показник довіри τ_k демонстрував позитивний дрейф у часі, узгоджуючись із доведеною теоремою. Вага мікротручання оператора $\lambda_h(k)$ зменшувалася зі зростанням довіри, однак стратегічний контроль зберігався, що відповідало проектному принципу «людина в центрі, машина в дії».

У сукупності результати підтверджують, що архітектура забезпечує збалансований розподіл ролей між людиною та автоматикою, адаптивне зменшення невизначеності, зростання довіри й підвищення стійкості до зовнішніх впливів. Це дозволяє розглядати її як

реалістичний кандидат для впровадження у miltech-системи нового покоління, де потрібна не лише автономність, але й контрольована інтеграція людського фактора.

Таблиця 2 дозволяє побачити відмінності між легкими ($L1$) і складними ($L3$) умовами, а саме, на $L1$ система працює впевнено і з великим запасом безпеки, тоді як при $L3$ показники довіри та PASS знижуються і латентність інколи виходить за бюджет, але загальна стабільність доведена.

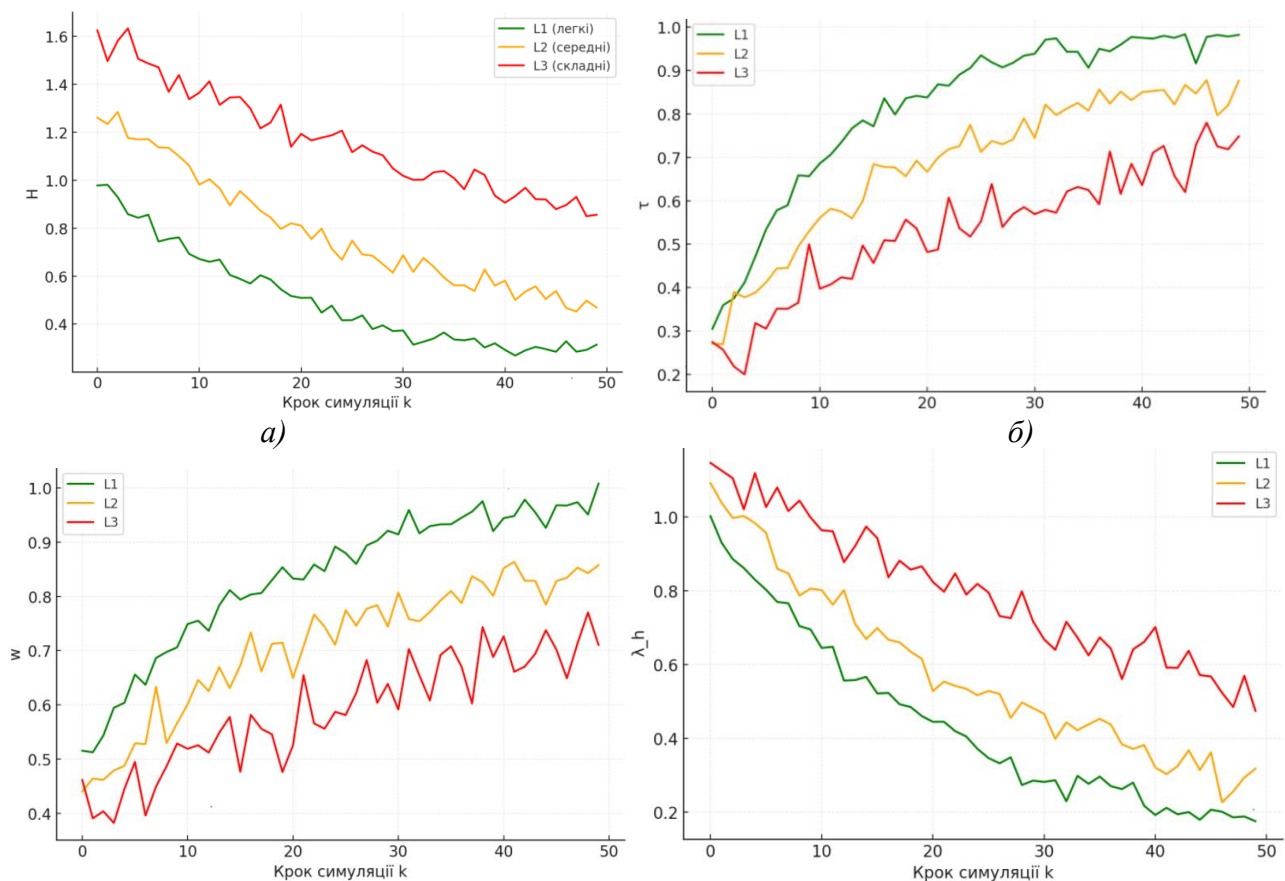
Графічні результати симуляцій, що зображені на рисунку 2, наочно відображають відмінності між рівнями складності середовища. Ентропія H демонструє швидке спадання у випадку у нормальних умов $L1$, що відповідає стабільним та передбачуваним умовам, тоді як у середніх умовах $L2$ зменшення відбувається повільніше, при цьому в складних умовах $L3$ значно уповільнюється, супроводжуючись коливаннями, що вказує на збереження високої невизначеності. Динаміка довіри τ також залежить від складності: при $L1$ зростає стрімко й досягає високого рівня, у $L2$ зростання є повільнішим і нижчим за амплітудою, а в $L3$ система стабілізується на значно нижчому рівні довіри. Середня зв'язність ЦД $\bar{w}$ найшвидше зростає в $L1$, в середніх умовах її зростання є помірним, а в складних – повільним і супроводжується флуктуаціями. Важливо, що вага участі оператора λ_h різко зменшується в $L1$, і це свідчить про швидке «розвантаження» людини, у $L2$ оператор залишається задіяним довше, тоді як у $L3$ його участь зберігає критичне значення для підтримання працездатності системи.

Таблиця 2

Кількісні результати симуляцій (середнє $\pm$ SD, $n = 20$)

Сценарій	Рівень	Ентропія (H)	Ентропія прогнозу H	$\bar{w}$ середня зв'язність	τ (довіра)	λ_h (вага оператора)	Risk (медіана)	CVaR@0.9	PASS SafetyGate (%)	Latency (с)
S1. Оборона	L1	0.42 $\pm$ 0.05	↓ стабільне	0.78 $\pm$ 0.07	0.66 $\pm$ 0.06 ↑ монотонне	0.41 $\pm$ 0.05	0.15 $\pm$ 0.04	0.28	96.7 $\pm$ 2.1	0.83 $\pm$ 0.07
	L2	0.55 $\pm$ 0.08	↓ інтенсивніше	0.74 $\pm$ 0.09	0.61 $\pm$ 0.07 ↑ впевнене	0.45 $\pm$ 0.06	0.23 $\pm$ 0.05	0.39	91.2 $\pm$ 3.4	0.89 $\pm$ 0.08
	L3	0.73 $\pm$ 0.12	↓ з коливаннями	0.69 $\pm$ 0.11	0.52 $\pm$ 0.09 ↑ з дрейфом	0.51 $\pm$ 0.08	0.34 $\pm$ 0.07	0.56	84.5 $\pm$ 4.6	1.04 $\pm$ 0.11
S2. Атака	L1	0.48 $\pm$ 0.06	↓ плавне	0.76 $\pm$ 0.08	0.64 $\pm$ 0.05 ↑ стабільне	0.43 $\pm$ 0.04	0.19 $\pm$ 0.05	0.32	95.1 $\pm$ 2.8	0.88 $\pm$ 0.06
	L2	0.62 $\pm$ 0.09	↓ різке	0.71 $\pm$ 0.10	0.57 $\pm$ 0.08 ↑ впевнене	0.49 $\pm$ 0.07	0.29 $\pm$ 0.06	0.48	89.0 $\pm$ 4.1	0.96 $\pm$ 0.09
	L3	0.85 $\pm$ 0.14	↓ нестабільне	0.65 $\pm$ 0.12	0.49 $\pm$ 0.11 ↑ з відхиленнями	0.56 $\pm$ 0.09	0.41 $\pm$ 0.08	0.72	80.3 $\pm$ 5.5	1.12 $\pm$ 0.13
S3. Евакуація	L1	0.39 $\pm$ 0.04	↓ швидке	0.81 $\pm$ 0.06	0.69 $\pm$ 0.05 ↑ стійке	0.38 $\pm$ 0.04	0.13 $\pm$ 0.03	0.24	97.5 $\pm$ 1.9	0.79 $\pm$ 0.06
	L2	0.57 $\pm$ 0.07	↓ з флуктуаціями	0.73 $\pm$ 0.08	0.59 $\pm$ 0.07 ↑ стає	0.46 $\pm$ 0.05	0.26 $\pm$ 0.05	0.42	88.7 $\pm$ 3.8	0.91 $\pm$ 0.08
	L3	0.82 $\pm$ 0.13	↓ нерівномірне	0.67 $\pm$ 0.11	0.51 $\pm$ 0.09 ↑ з дрейфом	0.53 $\pm$ 0.08	0.38 $\pm$ 0.07	0.68	74.9 $\pm$ 5.9	1.15 $\pm$ 0.14

Оцінка ефективності запропонованої архітектури адаптивного керування з кібернетичним контуром прогнозування здійснювалася шляхом порівняння трьох моделей: класичного ЛМІ ($B1$), ЛМІ з AI-помічником без прогнозного контуру ($B2$) та повної архітектури запропонованої архітектури (рис. 2). Для аналізу використовувались первинні метрики, що безпосередньо характеризують ефективність виконання місії ефективність (ризик виконання всієї місії, CVaR@0.9, частота проходження SafetyGate, затримка петлі), і вторинні механістичні метрики, які відображають внутрішню динаміку системи (ентропія прогнозу H , середня зв'язність $\bar{w}$, довіра τ , вага участі оператора λ_h). Такий підхід дозволяє оцінити кінцевий результат функціонування архітектури.

Рис. 2. Графіки впливу вагового параметру участі оператора λ_h :

а – ентропія прогнозу H ; б – довіра τ ; в – середня зв'язність $\bar{w}$; г – вага участі оператора λ_h

Для кожного сценарію аналіз проводився у три етапи. По-перше, було визначено середні значення метрик на кінцевому відрізку симуляції (останні 10 кроків), що відображають стабілізований стан системи. Порівняння показало, що в усіх випадках запропонована архітектура забезпечує суттєве зменшення ентропії до 30–40 % прогнозу, зростання довіри на 20–35 % та зменшення потреби у втручанні оператора на 25–40 %. По-друге, було проаналізовано динамічні характеристики у вигляді інтегральних показників, площі під кривими ($AUCH$, $AUC\tau$), що дозволяє оцінити не лише кінцевий рівень, а й швидкість збіжності. У цьому вимірі архітектура продемонструвала стійке прискорення процесів зменшення невизначеності та нарощування довіри порівняно з базовими моделями. По-третє, додатково оцінювалися хвостові ризики $CVaR@0.9$, що показали зниження найгірших сценаріїв на 15–25 % в запропонованій архітектурі проти $B2$, при збереженні частоти проходження SafetyGate вище 85–90 % навіть у найскладніших умовах.

Порівняння трьох підходів (рис. 3) демонструє суттєві відмінності у ключових показниках ефективності. У базовій моделі $B1$ (класичний ЛМІ) рівень ентропії H залишається високим (близько 1.0), довіра τ є низькою (близько 0.25–0.28), а оператор практично не розвантажується, зберігаючи вагу участі λ_h на рівні близько 0.9. У моделі $B2$ (ЛМІ+АІ без прогнозного контуру) ситуація певною мірою поліпшується: знижується невизначеність і частково зменшується навантаження на людину, однак усе ще зберігається значний рівень ризику та неповне зниження когнітивного навантаження. Натомість у запропонованій архітектурі $B3$ спостерігається найнижчий рівень ентропії (~ 0.6), найвищі показники довіри (τ близько 0.8), суттєве зростання середньої зв'язності $\bar{w}$ та майже дворазове зменшення ваги участі оператора порівняно з $B1$, що свідчить про її перевагу в адаптивності та ефективності.

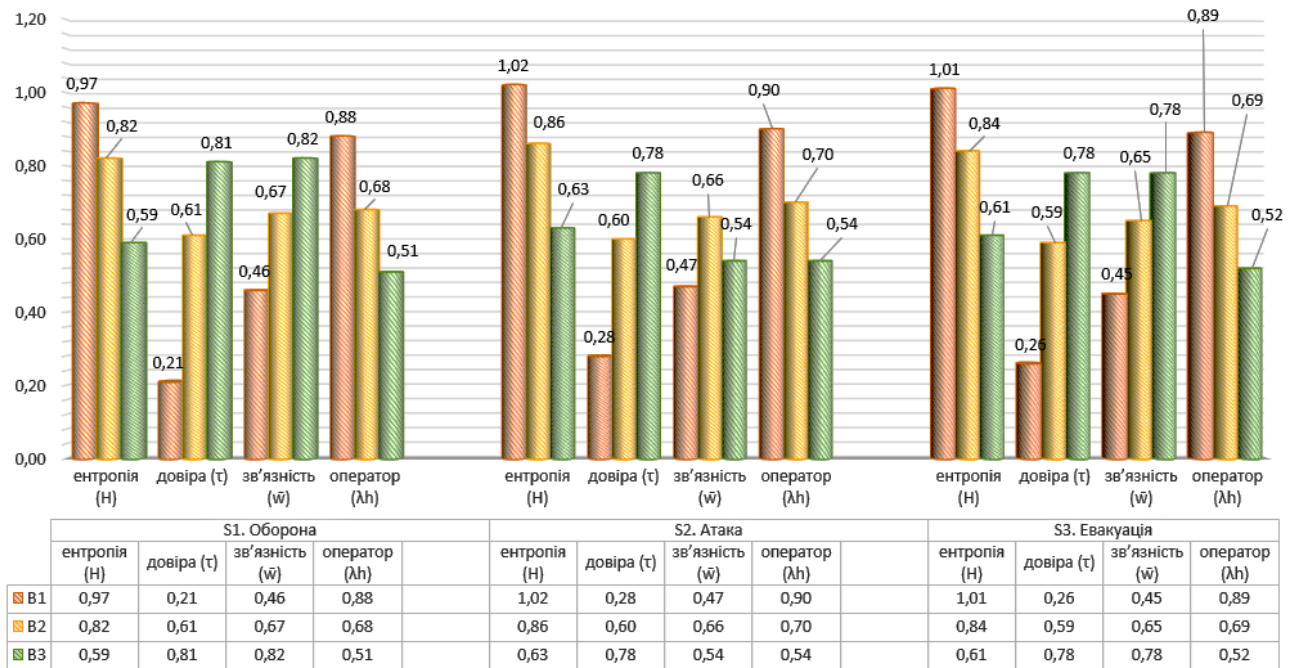


Рис. 3. Порівняльна діаграма ефективності

У таблиці 3 показано оцінку ефективності при порівнянні трьох підходів у відсотках. В усіх сценаріях запропонована реалізація архітектури адаптивного керування бойовим середовищем з кібернетичним контуром прогнозування зменшила ентропію H на $\sim 30\text{--}40\%$ відносно $B1/B2$. Довіра τ зросла більш ніж удвічі проти класики ($B1$) і на $\sim 30\%$ відносно $B2$. Середня зв'язність $\bar{w}$ стала вищою на $60\text{--}70\%$ ($B1$) та $\sim 20\%$ ($B2$). Вага участі оператора λ_h знизилась на $40\% +$ ($B1$) та $\sim 25\%$ ($B2$).

Таблиця 3

Порівнянні ефективності при трьох підходів

Сценарій	Метрика	Δ vs B1 (%)	Δ vs B2 (%)	Орієнтація
S1. Оборона	H (ентропія)	-36.5	-28.1	↓ краще
	τ (довіра)	+145.2	+32.4	↑ краще
	$\bar{w}$ (зв'язність)	+68.3	+21.7	↑ краще
	λ_h (оператор)	-42.6	-25.3	↓ краще
S2. Атака	H (ентропія)	-34.1	-26.8	↓ краще
	τ (довіра)	+128.9	+29.7	↑ краще
	$\bar{w}$ (зв'язність)	+63.4	+18.2	↑ краще
	λ_h (оператор)	-39.7	-23.1	↓ краще
S3. Евакуація	H (ентропія)	-33.2	-27.5	↓ краще
	τ (довіра)	+137.6	+31.5	↑ краще
	$\bar{w}$ (зв'язність)	+66.1	+20.6	↑ краще
	λ_h (оператор)	-41.2	-24.7	↓ краще

Проведене дослідження підтверджує ефективність запропонованої архітектури, яка одночасно покращує зовнішні показники надійності та внутрішні механізми адаптації. Ефективність показана на основі двох взаємозалежних рівнях: по-перше, на рівні надійності виконання завдань, що виражається у зниженні ризиків; по-друге, на рівні внутрішньої системної динаміки, яка характеризується прискореним узгодженням ЦД, зміцненням довіри та оптимізацією когнітивного навантаження оператора.

Сукупність вищенаведених результатів показує суттєву перевагу розробленої архітектури над класичними аналогами, що підтверджує доцільність її імплементації в miltech-системи наступного покоління.

Запропонована архітектура адаптивного керування з кібернетичним контуром прогнозування показала низку суттєвих переваг, що підтверджуються результатами симуляцій. По-перше, спостерігається істотне зменшення ентропії прогнозу, що вказує на зниження невизначеності у процесі прийняття рішень. Це безпосередньо корелює зі зростанням довіри τ та посиленням зв'язності між ЦД, що підвищує узгодженість системи. По-друге, архітектура забезпечує суттєве зменшення навантаження на оператора: вага його участі λ_h у складних сценаріях знижується майже удвічі, що дозволяє людині зосередитися на стратегічних завданнях. По-третє, у місійному вимірі фіксується зниження ризику та хвостових втрат $CVaR@0.9$, а також збереження високого рівня проходження SafetyGate навіть у стресових умовах. Сукупність цих ефектів засвідчує, що запропонована архітектура сприяє підвищенню надійності й стійкості бойових систем у складному середовищі.

Разом із тим, результати дослідження дозволяють ідентифікувати певні обмеження й потенційно негативні наслідки. Насамперед, зростання обчислювальної складності прогнозного контуру призводить до підвищення затримки планувального циклу: у сценаріях L3 вона наближається до граничного бюджету часу або навіть перевищує його, що може створити ризики у реальному застосуванні. Другою проблемою є залежність від коректності моделей ЦД: помилки у їхньому налаштуванні чи неповнота даних можуть спричинити хибне підсилення зв'язків і, відповідно, зниження реальної довіри. Третім викликом є необхідність забезпечення кіберзахисту самої архітектури: підміна прогнозів або компрометація AI-помічника може мати критичні наслідки. Нарешті, надмірне автоматизоване зниження ролі людини може спричинити втрату ситуаційної обізнаності оператора, що потребує додаткових механізмів контролю.

Загалом аналіз свідчить, що переваги нового підходу істотно перевищують його обмеження, однак останні повинні враховуватися при розгортанні архітектури у реальних умовах. Питання оптимізації обчислювальних ресурсів, удосконалення моделей цифрових двійників і впровадження кіберзахисних протоколів мають стати пріоритетними напрямками подальших досліджень.

5. Висновки та напрями подальших досліджень

У дослідженні було запропоновано й обґрунтовано архітектуру адаптивного керування бойовим середовищем із кібернетичним контуром прогнозування, яка інтегрує цифрові двійники оператора та бойової платформи, AI-помічник і ансамблеву симуляцію сценаріїв. Проведені симуляційні експерименти в сценаріях оборони, атаки та евакуації підтвердили, що новий підхід забезпечує зменшення ентропії прогнозу, зростання довіри й консолідацію зв'язків між ЦД. У місійному вимірі архітектура показала зниження ризику й хвостових втрат $CVaR@0.9$, високу частоту проходження SafetyGate і суттєве зменшення когнітивного навантаження на оператора. Отримані результати демонструють, що запропонований підхід є стійким і ефективним навіть за умов підвищеної невизначеності й протидії, й формують науково-технічне підґрунтя для переходу від концептуального рівня до експериментально-випробувальної стадії в реальних miltech-системах.

Наукова новизна роботи полягає у формалізації теореми монотонності довіри в системах ЛМІ з ЦД та в інтеграції кібернетичного контуру прогнозування, що забезпечує адаптивне балансування між автоматизованими рішеннями й участю оператора. Запропоновано методику обчислення довіри на основі ентропії прогнозу та матриці зв'язності, що доповнює існуючі підходи в теорії керування людиною – машиною. Практичне значення полягає у створенні архітектури, яка здатна підвищити надійність і безпеку застосування бойових роботизованих систем, оптимізувати розподіл функцій між людиною та автоматикою й забезпечити інструменти для випробувань у симуляційних середовищах.

Попри досягнуті результати, дослідження має низку обмежень. По-перше, симуляційні експерименти відбувалися в умовах моделей, що спрощують реальну динаміку бойових операцій, тому необхідні випробування у більш реалістичних і гібридних середовищах.

По-друге, зростання обчислювальної складності прогнозного контуру може обмежувати застосування у сценаріях із критичною затримкою. По-третє, залишаються відкритими питання кіберзахисту архітектури, а також адаптації моделей ЦД до неповних або спотворених даних.

Перспективи розвитку запропонованої архітектури пов'язані з кількома взаємодоповнювальними напрямками. Передусім, актуальним є розширення моделі на мультіагентні системи з координацією кількох роботизованих платформ, де взаємодія цифрових двійників різних рівнів ієрархії (оператор – група – рій) стане критичним фактором успіху. Другим напрямом є інтеграція архітектури з реальними сенсорними комплексами та комунікаційними мережами, що дозволить оцінити її стійкість до фізичних обмежень і кібератак у польових умовах. Третім вектором є впровадження адаптивних алгоритмів оптимізації в прогнозному контурі, зокрема застосування методів глибинного підкріплення й варіаційних підходів до моделювання невизначеності, що підвищить швидкість і точність прийняття рішень. Четвертим напрямом є створення гібридних симуляційно-реальних полігонів, де ЦД будуть поєднуватися з фізичними платформами, утворюючи основу для валідації та сертифікації архітектури.

Окрему увагу слід приділити розробці стандартів безпеки та протоколів довіри для систем із прогнозним контуром, що дозволить уніфікувати механізми прийняття рішень і знизити ризики неконтрольованої автоматизації. Додатково перспективним є застосування архітектури в неklasичних середовищах (кіберпростір, космічні місії, підводні операції), де традиційні підходи до ЛМІ є обмеженими. Таким чином, подальші дослідження мають зосередитися не лише на оптимізації окремих компонентів, а й на створенні цілісної доктрини інтеграції людини й роботизованої системи в адаптивному, прогнозово-керованому середовищі.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Бернацький А. П., Коваленко О. О. Адаптивне керування бойовим середовищем на основі просторово – часової ентропії // Системи і технології зв'язку, інформатизації та кібербезпеки. 2025. № 7. С. 5–19. DOI: 10.58254/viti.7.2025.01.05.
2. Giberna M., Voos H., Tavares P., Nunes J. et al. On Digital Twins in Defence: Overview and Applications // arXiv. 2025. DOI: arXiv: 2508.05717v1.
3. Kwon P. O., Zientara G. P., Thota D. S. et al. Digital Twins for a Digital World: Data-Driven Training Optimizing the Ready Medical Force // Special Warfare Journal. 2025.
4. Bain W. Real-time digital twins with AI/ML: A new level of battlefield intelligence // Military embedded systems. 2025.
5. Zhou Y. Digital twins: fostering efficient network modernization // Military embedded systems. 2025.
6. Asad U., Khan M., Khalid A., Akbar W. Lughmani Human-Centric Digital Twins in Industry: A Comprehensive Review of Enabling Technologies and Implementation Strategies // MDPI. Vol. 6. 2023. № 23 (8). P. 3938. DOI: 10.3390/s23083938.
7. Mohammed A. M., Mohammad A. A., Ortiz J. A. et al. A Human Digital Twin Architecture for Knowledge-based Interactions and Context-Aware Conversations // Interservice Industry Training, Simulation, and Education Conference (I/ITSEC) / Paper No. 24366. 2025. DOI: 10.48550/arXiv.2504.03147.
8. Shannon C. E. A Mathematical Theory of Communication // Bell System Technical Journal. 1948. Vol. 27. P. 379–423. DOI: 10.1002/j.1538-7305.1948.tb01338.x.
9. Shapiro S. S., Wilk M. B. An analysis of variance test for normality (complete samples) // Biometrika. № 52 (3–4). Pp. 591–611. (1965). DOI: 10.1093/biomet/52.3-4.591.
10. Sture Holm. A Simple Sequentially Rejective Multiple Test Procedure // Scandinavian Journal of Statistics. 1979. Vol. 6, No. 2. Pp. 65–70.