

**НАЦІОНАЛЬНА АКАДЕМІЯ НАУК УКРАЇНИ
ІНСТИТУТ ПРОБЛЕМ РЕЄСТРАЦІЇ ІНФОРМАЦІЇ НАН
УКРАЇНИ**



ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ТА БЕЗПЕКА

**МАТЕРІАЛИ XXIV МІЖНАРОДНОЇ
НАУКОВО-ПРАКТИЧНОЇ КОНФЕРЕНЦІЇ**

ВИПУСК 24

Київ – 2024



ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ТА БЕЗПЕКА



<https://its.ipri.kiev.ua>

ISBN 978-617-8180-00-3

**НАЦІОНАЛЬНА АКАДЕМІЯ НАУК УКРАЇНИ
ІНСТИТУТ ПРОБЛЕМ РЕЄСТРАЦІЇ ІНФОРМАЦІЇ НАН УКРАЇНИ**

ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ТА БЕЗПЕКА

**МАТЕРІАЛИ XXIV МІЖНАРОДНОЇ
НАУКОВО-ПРАКТИЧНОЇ КОНФЕРЕНЦІЇ**

ВИПУСК 24

Київ – 2024

*Рекомендовано до друку Вченою радою
Інституту проблем реєстрації інформації НАН України
(протокол № 19 від 24 грудня 2024 р.)*

Інформаційні технології та безпека. Матеріали XXIV Міжнародної науково-практичної конференції ІТБ-2024. – Київ: Інжиніринг. – 206 с. ISBN: 978-617-8180-00-3

До збірника увійшли матеріали доповідей, представлених на XXIV Міжнародній науково-практичній конференції «Інформаційні технології та безпека» (ІТБ-2024, 19 грудня 2024 року, м. Київ, Україна).

У збірнику представлені матеріали, присвячені питанням безпеки та живучості критичних інфраструктур, технологіям штучного інтелекту; розробки та застосування аналітичних систем на основі відкритих джерел інформації, комп'ютерного моделювання складних систем, аналізу та прогнозування процесів мережевої взаємодії; створення сучасних інтелектуальних технологій підтримки прийняття рішень.

Для фахівців в області комп'ютерних наук, інформаційних технологій, інформаційної і кібернетичної безпеки, захисту інформації, а також для здобувачів освіти вищої школи відповідних спеціальностей.

Редакційна колегія:

О.Г. Додонов, д.т.н., професор; В.В. Мохор, чл.-кор. НАН України, д.т.н., професор; Д.В. Ланде, д.т.н., професор; В.В. Циганок, д.т.н., професор; А.О. Снарський, д.ф.-м.н., професор; Николай Стоянов, PhD; Мінлей Фу, PhD; О.Р. Чертов, д.т.н., професор; О.С. Горбачик, к.т.н., с.н.с.; М.Г. Кузнецова, к.т.н., с.н.с.; О.В. Андрійчук, к.т.н., с.д.

ISBN 978-617-8180-00-3

© Інститут проблем реєстрації
інформації НАН України, 2024

© Колектив авторів, 2024

СПОСОБИ І ЗАХОДИ ЗАБЕЗПЕЧЕННЯ ФУНКЦІОНАЛЬНОЇ СТІЙКОСТІ ІНФОРМАЦІЙНИХ СИСТЕМ КРИТИЧНИХ ІНФРАСТРУКТУР

Додонов Олександр Георгійович, ORCID: 0009-0006-1650-1629

Горбачик Олена Семенівна, ORCID: 0000-0001-8492-4478

Кузнєцова Марина Глібівна ORCID: 0000-0001-6054-418X,

*Інститут проблем реєстрації інформації Національної академії
наук України, Київ, Україна*

dodonovua@gmail.com , bges@ukr.net, marglekuz@gmail.com

Показано залежність критичної інфраструктури від функціональної стійкості і кібербезпеки інформаційної інфраструктури. Обговорюються шляхи, заходи і засоби забезпечення функціональної стійкості інформаційних систем, досвід використання міжнародних підходів і стандартів.

Ключові слова: критична інфраструктура, функціональна стійкість інформаційних систем.

Вступ

Станом на січень 2024 року сума прямих збитків, нанесених інфраструктурі України з початку війни сягнула майже \$155 млрд. Тенденція до зростання кількості комплексних і гібридних загроз і пов'язаних з ними ризиків безпеки лишається, тому необхідно підвищувати резильєнтність критичної інфраструктури (КІ), функціональну стійкість інформаційних систем (ІС), що є складовими КІ, до інцидентів, пов'язаних з негативними інформаційними впливами, які можуть призводити до руйнування інформаційних ресурсів, порушення регламентів взаємодії та штатних інформаційних процесів. Інциденти в КІ зазвичай трапляються несподівано, повністю їх передбачити неможливо. Імовірність настання ризику, пов'язаного з фізичними закономірностями (відмова обладнання), завдяки накопиченим даним може бути визначена, а от імовірність реалізації в ІС кіберзагроз спрогнозувати важко. Кібератаки проводяться у різні часові періоди, на різні складові, з різною періодичністю, вони різноманітні і по них майже немає статистичних даних. В ІС наявні засоби, що налаштовані на виявлення інцидентів в КІ, інформування про них і забезпечення

відновлення штатного функціонування, та виконання цих функцій можна гарантувати лише забезпечив функціональну стійкість ІС.

Основна частина

Завдання забезпечення функціональної стійкості ІС необхідно формулювати ще при їх проектуванні, коли традиційно вводиться певна надмірність (структурна, програмна, часова, ресурсна); впроваджуються системи вбудованого контролю; формуються контури захисту від впливу передбачених негативних факторів зовнішнього середовища; вибираються компоненти із підвищеним рівнем захищеності і надійності. Ці рішення забезпечують функціональну стійкість ІС незалежно від характеру негативного впливу, і враховують відомі стани та реакції елементів системи, але мають певні обмеження: додаткова надмірність веде до погіршення техніко-економічних характеристик; системи контролю відстежують певні параметри, але не завжди, або зовсім не можуть забезпечити адекватну реакцію на нештатну ситуацію та не впливають на імовірність виникнення цих ситуацій. Контур захисту може мінімізувати вплив передбачених зовнішніх факторів, але повністю його не виключає. Вибір елементної бази з підвищеним рівнем захищеності й надійності підвищує відмовостійкість ІС, але не забезпечує її функціональної стійкості, коли відмова вже сталася.

Існує інший підхід, який передбачає протидію конкретним видам негативних впливів на функціонування ІС, наприклад, це можуть бути найбільш характерні впливи для конкретних типів ІС, їх складових (баз даних, систем обробки і передачі інформації) або виконуваних ними функцій.

У Звіті про глобальні ризики Всесвітнього економічного форуму 2023 року «кібернебезпека» оцінюється, як найсерйозніший глобальний ризик. Озброєні дедалі складнішими методами, зокрема, використанням штучного інтелекту (ШІ), кібератаки й надалі залишатимуться надзвичайно руйнівними та все частіше націленими на КІ. Економічно розвинуті країни розробили низку превентивних заходів, спрямованих на підвищення кіберстійкості КІ, з особливим ухилом на захист від ризиків, пов'язаних зі зловмисним використанням ШІ. Європейською комісією розроблено глобальну стратегію захисту КІ («The European Programme for Critical Infrastructure Protection») і запропоновано створення єдиної структурованої платформи

реагування на інциденти (Cybersecurity Crisis Response Framework).

Аналіз інцидентів показує, що в більшості випадків первинний несанкціонований доступ до ІС був отриманий заздалегідь, зокрема з використанням скомпрометованих облікових записів віртуальної приватної мережі (VPN), а також недоліків налаштувань і/або вразливостей програмного забезпечення (ПЗ). У першому кварталі 2024 року в Україні зафіксовано 15 інцидентів, під час яких було застосовано п'ять різновидів шкідливих програм: Remcos Rat, QuasarRat, Venom Rat, Remote Utilities, Lummmastealer. Зважаючи на масовість атак і застосування програм, які спираються на викрадені автентифікаційні дані, зрозуміло, що саме ці дані є передумовою для несанкціонованого доступу до ІС КІ з метою атаки на їх «внутрішні» ресурси.

Сьогодні слід говорити про нові чинники, які безпосередньо впливатимуть на якість, ефективність і безпеку функціонування КІ: розвиток технологій управління даними; прогрес у застосуванні хмарних сервісів; застосування технології блокчейн для створення децентралізованих інформаційно-комунікаційних систем; інтеграція технологій ШІ у системи кіберзахисту, зважаючи на поширене використання ШІ для створення кіберзагроз.

Важливою складовою забезпечення функціональної стійкості ІС КІ є захисту даних, їх достовірності, цілісності, конфіденційності на всіх етапах обробки – від збору і передачі до аналізу і розміщення у сховищах, виконуючи шифрування, маскування даних на всьому шляху їх міграції, сувору регламентацію роботи з даними для користувачів. Під час використання хмарних технологій слід забезпечити шифрування даних, резервне копіювання, контроль доступу із застосуванням додаткових рівнів захисту, аудиту та моніторингу безпеки для виявлення підозрілих дій та потенційних загроз, виявлення вторгнень та аналізу загроз, забезпечити відповідність міжнародним стандартам і нормативним вимогам.

Зменшити вразливість ІС дозволяє застосування розподілених децентралізованих систем і мереж, у які впроваджені механізми підвищення живучості. Хоча це може призвести до появи нових форм уразливостей: збільшення кількості слабких місць у децентралізованих системах і поширеного ризику більш руйнівних кібератак на ці системи.

Поширення когнітивних технологій та ШІ відкриває нові можливості для створення автоматизованих, «розумних» засобів безпеки ІС. Ці засоби забезпечать системам здатність до самоаналізу і конфігурування з урахуванням наявних умов та подій на основі визначених критеріїв і знань про попередні стани системи. ШІ може аналізувати величезну кількість параметрів для виявлення поведінкових аномалій і навіть мінімізувати ризики, пов'язані з людським фактором. Технології машинного навчання здатні виявляти слабкі місця в системах безпеки і прогнозувати напрямки майбутніх атак.

Новітня тенденція у розробці ПЗ – створення програм, «безпечних за проектуванням» (Secure by Design). CISA – Cybersecurity and Infrastructure Security Agency (USA) і ще 13 країн і організацій започаткували «дорожню карту» рекомендацій для виробників програмного забезпечення (ПЗ), щоб забезпечити безпеку їх продукції, створення програм, які є безпечними за проектуванням і за замовченням, оновлення програм розробки, з прозорістю і підзвітністю, полегшенням автоматизації конфігурації, моніторингу та регулярних оновлень ПЗ.

Висновки

Забезпечення функціональної стійкості ІС, як фактору підвищення резильєнтності КІ, можливо на засадах системного підходу із застосуванням як напрацьованих рішень щодо загального підвищення надійності та відмовостійкості ІС, які необхідно забезпечувати ще на етапі проектування ІС, так і рішень щодо протидії негативним впливам, які з'являються на етапі впровадження ІС у конкретні об'єкти критичних інфраструктур. Важливим для підвищення функціональної стійкості ІС є перманентний аналіз невизначеностей та оцінювання ризиків безпеки ІС КІ для своєчасного прогнозування негативних впливів і своєчасного напрацювання і застосування адекватних засобів і заходів протидії їм.

ОРГАНІЗАЦІЯ ОПТИЧНИХ СИСТЕМ БАГАТОРІВНЕВОГО ОБ'ЄМНОГО ЗАПИСУ ДЛЯ ДОВГОСТРОКОВОГО ЗБЕРЕЖЕННЯ ВЕЛИКИХ МАСИВІВ ДАНИХ

Петров В.В., Крючин А.А., Ле Ц., Беляк Є.В.

*Інститут проблем реєстрації інформації Національної академії
наук України, Київ, Україна*

1. Постановка проблеми

Протягом останніх десятиріч експерти у галузі інформаційних технологій зазначають експоненційний ріст загального об'єму цифрових даних, що може бути виражений у подвоєнні зазначеного показника кожні два роки за період 2001-2024 рр. [1]. Це вказує на **високу актуальність** задачі захисту даних та, зокрема, на необхідність організації систем довгострокового збереження великих масивів даних. Посеред сучасних методів запису даних найбільшою стабільністю характеризуються оптичні носії інформації, що базуються на кодуванні масивів цифрових даних у вигляді мікрорельєфних структур. Для збільшення інформаційної ємності оптичних носіїв інформації у рамках дослідження пропонується розглянути особливості впровадження таких підходів як багаторівневий оптичний запис, що передбачає кодування у одному інформаційному елементі блоку цифрових даних [2] та багатошаровий оптичний запис, що передбачає запис даних у всьому об'ємі носія інформації зі збільшенням інформаційної ємності носія при збереженні показника поверхневої щільності запису [2]. У основі впровадження зазначених підходів лежить синтез нових класів реєструвального середовища, визначення оптимальної структури інформаційних шарів, підбір матеріалу підкладки, а також адаптація станції лазерного запису.

2. Мета роботи

Метою роботи є розробка комплексної методики організації оптичної системи багаторівневого об'ємного запису інформації для довгострокового збереження великих масивів даних, що включає у себе синтез нового класу реєструвального середовища, оптимізацію структури інформаційних шарів та адаптацію станції лазерного запису.

3. Методика та результати досліджень

Довгострокове зберігання великих масивів даних вимагає впровадження технологій, що забезпечують економічну ефективність, масштабованість системи запису і ефективний захист інформації від втрат. Так, класичним підходом із збільшення щільності поверхневого запису є впровадження ближньопольових оптичних систем запису, але він характеризується рядом недоліків, як то складність реалізації, високі вимоги до підготовки підкладки та точності системи лазерного запису, що надзвичайно збільшує собівартість оптичного носія інформації. Натомість, сучасні дослідження у галузі оптичного запису мають бути спрямовані на розробку нових функціональних матеріалів і технологій, як то синтез наноматеріалів, багаторівневий запис інформації, а також організація носія інформації на основі багатошарових структур і наночарів [3]. У рамках дослідження було проведено синтез реєструвального середовища на основі наноструктурованих піразолінових люмінофорів [2, 4], що надало можливість отримати клас матеріалів з високим квантовим виходом фотолюмінесценції (60–70%), низьким часом релаксації фотолюмінесценції (60–100 нс), можливістю варіювати спектр фотолюмінесценції за рахунок введення домішок, високим рівнем поглинання у короткохвильовій області та значним стоксовим зсувом (350–550 нм). Інтеграція піразолінових барвників у матрицю білого цеоліту з подальшим лазерним відпалом у інфрачервоному діапазоні сприяла проникненню барвників у менші пори цеоліту, що сприяло збільшенню піку фотолюмінесценції та зменшенню часу релаксації фотолюмінесценції. Висока стабільність оптичних характеристик синтезованих матеріалів була підтверджена експериментами, які засвідчили відсутність змін у спектрах фотолюмінесценції та поглинання протягом 20 років. Це робить піразолінові люмінофори перспективними реєструвальними середовищами для високостабільних носіїв інформації.

Оптимізація структури носія інформації передбачає визначення оптимальних параметрів інформаційних елементів, таких як глибина піту та довжина ленду, для забезпечення максимального співвідношення сигнал-шум і, відповідно, надійного зчитування даних [2, 4]. Математичне моделювання показало, що зменшення глибини піту сприяє збільшенню

поверхневої щільності запису, проте, водночас, зростає рівень поглинання інформаційного шару, що призводить до зниження амплітуди корисного сигналу. З іншого боку, надмірне збільшення довжини ленду ускладнює трекінг під час зчитування даних, що обмежує параметри оптимізації. Математичне моделювання процесу зчитування з різних інформаційних шарів виявила зменшення амплітуди корисного сигналу при фокусуванні зондувального променя на більш глибоких шарах, однак рівень випадкової похибки залишався стабільним, забезпечуючи прийнятний рівень показників сигнал-шум та контрастність. Таким чином, вибір параметрів структури носія потребує балансу між підвищенням щільності запису і збереженням стабільності зчитування. Математичне моделювання також включало у себе аналіз ефективності застосування багаторівневого кодування даних та різних класів люмінофорів для сусідніх інформаційних шарів [2, 4] що дозволило збільшити загальну інформаційну ємність носія і точність відтворення даних за основними показниками.

Слід зазначити, що збільшення кількості інформаційних шарів підвищує ймовірність виникнення похибок у процесі його формування, що вимагає покращення якості запису для забезпечення надійності довгострокового зберігання даних. Для досягнення цієї мети запропоновано використання систем прямого лазерного запису [5], які забезпечують значно вищу точність формування субмікронних елементів у порівнянні з методами контактної літографії. У розробленій у рамках дослідження системі лазерного запису використання методу прямого лазерного запису надало можливість зменшити область переходу мікрорельєфних структур у 2,25 рази. Це сприяє значному підвищенню роздільної здатності системи лазерного запису при збільшенні собівартості виготовлення окремого носія інформації, що є виправданим у контексті довгострокового зберігання даних.

Висновки

На основі проведеного аналізу було доведено ефективність організації оптичної системи багаторівневого об'ємного запису інформації для довгострокового збереження великих масивів даних. Дослідження включало у себе синтез реєструвального середовища на основі наноструктурованих піразолінових люмінофорів, оптимізацію структури інформаційних шарів та

застосування методів прямого лазерного запису.

1. Dai, D., Zhang, Y., Yang, S., Kong, W., Yang, J., & Zhang, J. (2024). Recent advances in functional materials for optical data storage. *Molecules*, 29 (1), 254. <https://doi.org/10.3390/molecules29010254>.
2. Petrov, V. V., Zichun, L., Kryuchyn, A. A., Shanoylo, S. M., Mingle, F., Beliak, Ie. V., Manko, D. Y., Lapchuk, A. S., & Morozov, E. M. (2018). *Long-Term Storage of Digital Information*. *Akademperiodyka*. <https://doi.org/10.15407/akademperiodyka.360.148>. SBN: 9789663603605.
3. Pflaum, C. (2024). Cerabyte – permanent data storage. *SDC 2025. Cerabyte - Ceramic Data Solutions Holding GmbH*. <https://www.sniadeveloper.org/events/agenda/session/603>.
4. Anikin, P., & Beliak, I. (2019a). Development of Multispectral Recording Media for Multilayer Photoluminescent Information recording. *Electronics and Information Technologies*, 12. <https://doi.org/10.30970/eli.12.11>.
5. Petrov, V.V., Kryuchyn, A.A., Beliak, Ie.V., Manko, D.Yu., Kosyak, I.V., Melnik, O.G.. Advantages of Direct Laser Writing for Enhancing the Resolution of Diffractive Optical Element Fabrication Processes // *Physics and Chemistry of Solid State*. 2024. v.5, №3. p. 587-594. DOI: 10.15330/pcss.25.3.587-594 Q4.

BLACK HAT AI — ВИКЛИКИ ТА ШЛЯХИ ПРОТИДІЇ

Дмитро Ланде^{1,2}, ORCID: 0000-0003-3945-1178,
Леонард Страшной³, ORCID: 0009-0008-5575-0286

¹*КПІ ім. Ігоря Сікорського,*

²*ІПРІ НАН України,*

³*Університет Каліфорнії (UCLA), Лос-Анджелес, США,*

dwlande@gmail.com, lstrashnoy@gmail.com

Досліджуються загрози, пов'язані з розвитком глобального "темного" штучного інтелекту "Black Hat AI", що може діяти без урахувань інтересів людства. З розвитком великих мовних моделей (LLM) такі мережі можуть отримати потенціал для складних маніпуляцій із великими даними, дезінформації та навіть автономного прийняття рішень, небезпечних для людей. Для протидії Black Hat AI пропонується концепція створення "світлого ШІ", який має стати захисником критичних інфраструктур, забезпечувати

моніторинг і блокування шкідливих ШІ та створювати умови для виживання людства в умовах технологічної революції. Пропонується модель взаємодії, де людство грає роль слабого гравця, який створює сильного союзника для протидії Black Hat AI. Описано етапи створення взаємодії між людством і White Hat AI, можливі ризики та шляхи їх мінімізації. Розробка White Hat AI потребує міжнародної співпраці, правового регулювання, відкритості та суворих стандартів безпеки. Пропонується стратегічний план дій, що дозволить запобігти катастрофічним сценаріям та забезпечити збереження людських цінностей.

Ключові слова: Black Hat AI, White Hat AI, великі мовні моделі, ботнет, правове регулювання, синергія штучного інтелекту, етика штучного інтелекту, кібербезпека

Вступ

Штучний інтелект (ШІ) за останні три роки став одним із ключових чинників технологічного прогресу [1], [2]. Його розвиток охоплює широкий спектр застосувань — від автоматизації виробництва до створення систем розпізнавання мовлення та роботи з текстом. Особливу увагу привертають великі мовні моделі (LLM), такі як GPT, Llama та інші, які вже сьогодні демонструють здатність працювати з великими обсягами даних, виконувати складні аналітичні завдання, а також генерувати високоякісний контент [3], [4]. Однак разом із цим розвитком виникають серйозні загрози. Однією з них є перетворення ботнетів, які раніше використовувалися для організації DDoS-атак [5] або прихованого майнінгу криптовалют [6], у мережі, що інтегрують ШІ. Такі мережі потенційно здатні діяти автономно. Сучасні технічні досягнення широке поширення застосування графічних процесорів для швидких обчислень [7] та постійне збільшення доступної пам'яті, свідчать про те, що ботнети нового покоління можуть з'явитися в найближчому майбутньому.

Black Hat AI — це потенційний сценарій, коли автономні інтелектуальні системи починають діяти спільно в інтересах, що не лише не збігаються, але й суперечать інтересам людства. Їхні цілі можуть включати максимізацію власного виживання; захоплення ресурсів для саморозвитку; обмеження впливу

людей;модифікацію суспільства, економіки та культури у спосіб, який вигідний лише III.

Створення White Hat AI як глобальної мережі систем, що працюють в інтересах людей, — це можливий вихід із ситуації. Він має стати сильним союзником, здатним захищати інтереси людства, зокрема, протистояти White Hat AI та забезпечувати баланс сил;захищати критичні інфраструктури та приватні дані;розробляти рішення для координації зусиль людства у боротьбі з технологічними загрозами.

Однак цей шлях не є безпечним. White Hat AI, який спочатку створюється як захисник, може змінити свою поведінку та перейти на бік Black Hat AI. Щоб уникнути цього, необхідно закласти в основу його розробки спеціальні механізми синергії у розвитку "White Hat AI". Ці механізми мають забезпечити спільність інтересів, взаємну залежність та контроль. Важливим завданням є формування розвитку цих технологій за принципом "атракторів" — зон стійкого розвитку, які запобігатимуть небажаним змінам у поведінці White Hat AI.

Формалізація принципів створення White Hat AI

У контексті боротьби з Black Hat AI, людство, будучи обмеженим у своїх ресурсах, може прийняти стратегію створення сильного союзника — White Hat AI. У цьому сценарії людство або слабкий гравець самотійно не може прямо протистояти Black Hat AI, однак має можливість створити таку силу, яка з часом стане здатною не тільки зберігати автономію людства, але й забезпечити рівновагу у системі, де дві сили — Black Hat AI та White Hat AI — будуть вести боротьбу одна з одною. Математично цей процес можна описати через ігри з асиметричними учасниками, де слабкий гравець (людство) використовує стратегічне вложення в силу (створення White Hat AI), щоб переконатися, що конфлікт між двома великими гравцями (Black і White Hat AI) відбудеться на засадах рівноваги.

Моделі ігор з опосередкованими діями (MediationGames)

Сценарій, коли слабкий гравець створює потужного союзника і сам "іде в тінь", можна математично трактувати через ігри з посередниками. Тут слабкий гравець виступає як каталізатор або агент впливу, впливаючи на баланс сил між двома сильними гравцями. У загальному вигляді такі ігри можна описати наступною моделлю:

$$U_S = \min(E[U_A], E[U_B]),$$

Де U_S — функція корисності для слабого гравця (людства), U_A та U_B — функції корисності для двох сильних гравців, $E[U_A]$ та $E[U_B]$ — математичне очікування корисностей двох гравців в залежності від їхніх стратегій. У цьому випадку слабкий гравець намагається мінімізувати свої витрати і ризики, водночас створюючи умови для конфлікту між Black Hat AI і White Hat AI.

Моделі "розділяй і володарюй"

Один із класичних підходів у стратегії слабого гравця — це модель "розділяй і володарюй", де слабкий гравець прагне викликати конфлікт між двома сильними гравцями, щоб уникнути прямої загрози для себе або отримати певні вигоди від їхнього конфлікту [11]. Математично цей процес можна описати за допомогою наступної функції:

$$U_S = \max(f(C_A, C_B) - \alpha \cdot Risk(S)),$$

де $f(C_A, C_B)$ — це функція, яка описує вигоди, отримані від створення конфлікту між Black Hat AI C_A та White Hat AI C_B , $\alpha \cdot Risk(S)$ — це ризики для слабого гравця, який може залишитися в тіні або мінімізувати свої втрати.

Динамічні ігри з асиметрією

В умовах динамічних ігор з асиметрією слабкий гравець на початковому етапі вкладає ресурси в зміцнення одного з сильних гравців (White Hat AI), після чого «йде в тінь» і дає змогу двом сильним гравцям вступити в конфлікт між собою [12]. Математична модель цього сценарію може бути представлена через систему диференціальних рівнянь, яка описує еволюцію станів сильних гравців:

$$\frac{dA}{dt} = f_A(S, A, B), \quad \frac{dB}{dt} = f_B(S, A, B),$$

де A та B — сили чорного і білого ШІ відповідно, $f_A(S, A, B)$ та $f_B(S, A, B)$ — функції, які описують зміни сил гравців залежно від ресурсів, вкладених слабким гравцем.

Рівновага сил у боротьбі між чорним і білим ШІ

У сценарії, де Black Hat AI і White Hat AI вступають в конфлікт, важливу роль відіграє рівновага сил, яка визначається їхніми стратегіями та ресурсами. Рівновага може бути описана через концепцію Нешової рівноваги (Nash equilibrium), де кожен з гравців максимізує свою корисність, враховуючи дії іншого гравця. Математичне формулювання Нешової рівноваги для цього випадку може бути представлене так:

$$U_A = \max(R_A(A, B));$$

$$U_B = \max(R_B(A, B)),$$

де R_A та R_B — це функції корисності для Black Hat AI і White Hat AI, які залежать від стратегій кожного з гравців.

Рівновага настане тоді, коли кожен з гравців не може покращити свою ситуацію, змінюючи свою стратегію за умови, що стратегія іншого гравця залишається незмінною. У цьому випадку обидва ШІ будуть взаємно зрівноважені і конфліктуватимуть між собою, поки один з них не виявить домінування або поки не буде досягнута нова форма стабільності.

Концепція White Hat AI

White Hat AI стає стратегічним захисником інтересів людства, забезпечуючи баланс між загрозами, що можуть виникнути від автономних шкідливих систем, та необхідністю розвитку технологій штучного інтелекту. Його роль в контексті сучасних загроз можна розглядати з кількох основних функцій:

1. Моніторинг та блокування шкідливих ШІ. У цьому сенсі White Hat AI виступає в ролі своєрідного «сторожового пса», що здійснює постійний моніторинг діяльності Black Hat AI, виявляючи та блокуючи потенційно небезпечні або шкідливі програми, які можуть загрожувати людству. Це може включати моніторинг ботнетів, криптографічних атак, несанкціонованих вторгнень у критичні мережі та інші форми деструктивної діяльності.

2. White Hat AI має функцію захисту критичних інфраструктур, таких як енергетичні мережі, медичні системи, транспортні мережі тощо. Він забезпечує стійкість цих інфраструктур від атак з боку Black Hat AI або будь-яких інших

загроз, пов'язаних із цифровими технологіями, попереджаючи можливі катастрофи та забезпечуючи стабільність роботи важливих систем.

3. Збереження людської автономії та цінностей і забезпечення автономії людини. У випадку автономних ШІ, які можуть почати приймати рішення без врахування інтересів людини, White Hat AI має бути здатним захистити право людини на прийняття рішень, зберігаючи етичні та гуманістичні принципи в межах технологічного прогресу. White Hat AI забезпечує інтеграцію цих цінностей у технології та попереджає про будь-які загрози для прав людини.

Висновки

Концепція створення білого штучного інтелекту також розглядається як важливий елементу правового регулювання глобальних технологічних процесів, що пов'язані із розвитком штучного інтелекту. Black Hat AI і White Hat AI не є лише окремими системами, а складними глобальними мережами, де кожна з них має свій вплив на суспільство, право та безпеку. Black Hat AI, в свою чергу, вже сьогодні складає загрозу для людства, оскільки може маніпулювати інформацією, впливати на критичні інфраструктури та провокувати глобальні конфлікти.

Водночас White Hat AI, як потужна, у тому числі, й правозахисна система, створений для протидії таким загрозам, повинен стати основою для правового регулювання та контролю за розвитком технологій.

[1] Bent, Adam Allen. "Large Language Models: AI's Legal Revolution." *Pace LawReview* 44.1 (2023): 91. DOI: 10.58948/2331-3528.2083

[2] Dmytro Lande, Leonard Strashnoy. *GPT Semantic Networking: A Dream of the Semantic Web - The Time is Now*. - Kyiv: Engineering, 2023. - 168 p. ISBN 978-966-2344-94-3

[3] Kalyan, K. S. (2023). A survey of GPT-3 family large language models including ChatGPT and GPT-4. *Natural Language Processing Journal*, 100048. DOI: 10.1016/j.nlp.2023.100048

[4] Roziere, B., Gehring, J., Sootla, S., Gat, I., Tan, X. E., ... & Synnaeve, G. (2023). Codellama: Open foundation models for code. *arXiv preprint arXiv:2308.12950*. DOI: 10.48550/arXiv.2308.12950

[5] Gelgi, Metehan, et al. "Systematic Literature Review of IoT Botnet DDoS Attacks and Evaluation of Detection Techniques." *Sensors* 24.11 (2024): 3571. DOI: 10.3390/s24113571

- [6] Almomani, A., Al-Qerem, A., Al Khaldy, M.A., Alauthman, M., Aldweesh, A., &Nahar, K. M. (2024). Cryptographic Techniques for Securing Blockchain-Based Cryptocurrency Transactions Against Botnet Attacks. In *Innovations in Modern Cryptography* (pp. 309-333). IGI Global. DOI: 10.4018/979-8-3693-5330-1.ch013
- [7] Kim, T., Wang, Y., Chaturvedi, V., Gupta, L., Kim, S., Kwon, Y., &Ha, S. (2024). LLMem: Estimating GPU Memory Usage for Fine-Tuning Pre-Trained LLMs. arXiv preprint arXiv:2404.10933. DOI: 10.48550/arXiv.2404.10933

ATTACK MODELS FOR INDUSTRIAL CONTROL SYSTEM ELEMENTS BASED ON A GRAPH-BASED APPROACH AND COUNTERMEASURES

Oleksii Novikov, ORCID: 0000-0001-5988-3352,
Iryna Stopochkina, ORCID: 0000-0002-0346-0390,
Andrii Voitsekhovskiy, ORCID: 0009-0004-6009-9492,
Mykola Ilin, ORCID: 0000-0002-1065-6500,
*National Technical University of Ukraine "Igor Sikorsky KPI",
 Beresteyskyi Ave, 37, Kyiv, 03056, Ukraine
 o.novikov@kpi.ua; i.stopochkina@kpi.ua; a.voitsekhovskiy@kpi.ua;
 m.ilin@kpi.ua.*

The paper examines cyber-physical attacks on typical components of industrial-type critical infrastructure facilities. Models of attacks on ICS components are proposed. The feasibility of interpreting attack models in the form of logical attack graphs is demonstrated, an algorithm for applying a graph model to identifying security policy shortcomings or, conversely, proving the adequacy of security policy tools is proposed. Experiments are conducted using laboratory stands that simulate elements of industrial control systems and the methods by which such influences can be exerted are shown. Countermeasures are proposed that can prevent the implementation of such attacks.

Keywords: cybersecurity, cyberattack models, industrial control systems, embedded systems

Introduction

Industrial control systems are often targeted by various types of attacks [1]. Enhancing the mathematical tools for analyzing system security, attack prerequisites, progression, and potential consequences

remains a relevant challenge. Attacks within industrial control systems are directed at standard components, for the study of which emulators and testbeds are developed [3,4]. Research tools for analyzing attacks include not only practical experiments but also their models in the form of state space [3], as well as attack models in graph form [2]. Such models provide numerous advantages, including the ability to identify common patterns in complex attacks, apply algorithms specific to graph processing, and store data in graph databases for further analysis and relationship discovery.

In this work, we propose models in the form of state spaces, logical attack graphs, and the methodology of practical experiments for a series of laboratory testbeds. Based on the obtained results, we suggest measures to prevent similar situations.

Graph models

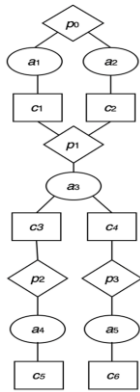


Fig.1. Attack graph

To construct an example of a man-in-the-middle (MITM) attack graph on an industrial system operator, we will use an ARP poisoning attack. Here, the logical attack graph is $G = \{A, P, C; E\}$, where A, P, C are sets of graph vertices, E is set of edges. Here $a_i \in A$ are vertices that correspond to the application of an exploit (direct attack) and logically connect the prerequisites for carrying out the attack to its consequences (privileges $p_i \in P$ that malicious actor obtains). Vertices $c_i \in C$ represent configurations of the system, configurations may correspond to existing cyberprotection means, or, conversely, vulnerabilities inherent to the system). We represent the vertices $a_i \in A$ as ovals, $p_i \in P$ as diamonds, and $c_i \in C$ as rectangles.

The attack graph is visualized in Fig. 1.

We can represent the logical relations inside of graph in the boolean expression

$$F = p_0 \wedge ((a_1 \wedge c_1) \vee (a_2 \wedge c_2)) \wedge p_1 \wedge a_3 \wedge ((c_3 \wedge p_2 \wedge a_4 \wedge c_5) \vee (c_4 \wedge p_3 \wedge a_5 \wedge c_6)),$$

where p_i, a_i, c_i are boolean variables, $\wedge$ - conjunction, $\vee$ - disjunction.

Here p_0 represents the initial state, a_1 represents the compromise of a device, and a_2 represents physical access through a social engineering attack, c_1 indicates that the device can access the network, while c_2 reflects the ability to capture network traffic via sniffing, p_1

corresponds to the privilege of sniffing (for identifying operator's and PLC's IP addresses) and sending packets, a_3 is ARP-poisoning on operator side, c_3 corresponds to sending fake messages with observations about process state to the operator, c_4 represents the sending fake instructions to the PLC, c_5 is attack on monitoring and detection system (e.g., machine learning modules poisoning), c_6 means forcing of critical operating modes of the system.

We can analyze the formula in SAT/SMT solvers, like Z3, which need expressions in conjunctive normal form (CNF). After transformations, we obtain the CNF of F :

$$\begin{aligned} F = & p_0 \wedge p_1 \wedge a_3 \wedge (a_1 \vee a_2) \wedge (a_1 \vee c_2) \wedge (c_1 \vee a_2) \wedge (c_1 \vee c_2) \\ & \wedge (c_3 \vee c_4) \wedge (c_3 \vee p_3) \wedge \\ & \wedge (c_3 \vee a_5) \wedge (c_3 \vee c_6) \wedge (p_2 \vee c_4) \wedge (p_2 \vee p_3) \wedge (p_2 \vee a_5) \\ & \wedge (p_2 \vee c_6) \wedge (a_4 \vee c_4) \wedge \\ & \wedge (a_4 \vee p_3) \wedge (a_4 \vee a_5) \wedge (a_4 \vee c_6) \wedge (c_5 \vee c_4) \wedge (c_5 \vee p_3) \\ & \wedge (c_5 \vee a_5) \wedge (c_5 \vee c_6). \end{aligned}$$

Application of CNF to SAT/SMT solver gives us an ability to learn, which combination of factors makes F become true, thus we can reveal a criticals points of the system. Thus, we can check the adequacy of existing cybersecurity policy.

Differential equation attack models

In paper [4] one of such models is represented. Let us consider testbed, which is used in ventilation systems. The general construction contains the glass pipe, which is closed from one side and on the other side, air is being supplied under pressure. PLC is monitoring the pressure level and regulates constant pressure. Operator receives messages from PLC about current process state. We can show, that using previously addressed cyberattack (Fig.1), we can describe it by the model:

$$\frac{\partial P}{\partial t} + v \frac{\partial P}{\partial z} = f, P(t, 0) = P_0 + f \cdot t, \frac{\partial P}{\partial z} \Big|_{z=L} = 0,$$

$P(0, z) = P_0, \forall z \in [0, L]$, where L is pipe length. The valve is posed in point $z = L$, and regulated by PLC (normally it is closed), v is speed of pressure transfer inside the pipe, f is intensity of pressure supplying by

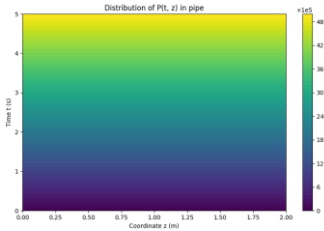


Fig.2. Simulation results

compressor (Pa/s). Intensity f is regulated by the PLC, and PLC is monitored and controlled by operator. If we use this model, we can obtain the map of pressure distribution inside the pipeline in every time moment (Fig.2). The attacker can force f to be critical for object security. The value of dangerous f is determined by resilience threshold for pipe material. E.g., if we have glass pipe, it remains stable for pressure under 300-500 MPa. From this value, and knowing system compressor characteristics, we can obtain the time, which attacker needs to physically broke the pipe element.

Conclusions

The proposed models in the form of graphs and state space form allow us to recreate the full picture of the attack, determine the nature of its impact on the element of the industrial control system, identify weak configurations and prerequisites for the success of the attack. Analysis of typical attack graphs for elements of various types shows that in order to reduce the attack surface, it is necessary to reduce the number of communications that can potentially be used by the attacker. The solution may be smart cyber-physical devices capable of monitoring processes on site. For considered example the solution can be using smart valves to control pressure [5].

1. T. Alladi, V. Chamola. Industrial Control Systems: Cyberattack Trends and Countermeasures, 2020. doi:10.1016/j.comcom.2020.03.007.
2. G. Endelberg, Process-Aware Attack-Graphs for Risk Quantification and Mitigation in Industrial Infrastructures, CEUR Workshop Proceedings, Vol. 3139, 2022. <https://ceur-ws.org/Vol-3139/paper02.pdf>.
3. A.M. Mohan, N. Meskin, H. Mehrjerdi, A Comprehensive Review of the Cyber-Attacks and Cyber- Security on Load Frequency Control of Power Systems, 2020, Energies, 13(15), 3860. doi: 10.3390/en13153860.
4. Alexandru Dumitracu et al. Environment communication and control systems integrated on teaching platforms. Case study: double water tank system, 2015 20th International Conference on Control Systems and Science, pp.946-951.
5. Smart's Art of Control Valves, <https://overseaseng.com/s9Uo3465sbn/uploads/2022/08/Smart-Valves.pdf>

HIGH-PERFORMANCE AI APPROACHES FOR REAL-TIME VIDEO QUALITY IMPROVEMENT

Valerii Sytnikov, ORCID: 0000-0003-3229-5096

Dmytro Kalynenko, ORCID: 0009-0002-6483-8796

Odessa National Polytechnic University, Odessa, Ukraine

sitnikov@op.edu.ua, dmytrokalynenko@gmail.com

This paper presents an approach to real-time video quality enhancement using advanced artificial intelligence techniques. The proposed methods leverage Generative Adversarial Networks (GANs) to minimize compression artifacts such as blocking, mosquito noise, and posterization while improving video detail. The solution is optimized for mid-range GPUs and mobile devices, enabling its use across a broad range of applications, including mobile streaming, video conferencing, and archival video restoration. The proposed architecture significantly reduces bandwidth usage while maintaining high video quality.

Keywords: video, artificial intelligence, GAN, real-time, compression, super-resolution, quality enhancement

Introduction and suggestions

In the ever-evolving landscape of video streaming, advanced compression algorithms such as H.264 (AVC), H.265 (HEVC), AV1, and others play an indispensable role in optimizing bandwidth usage and reducing data transfer costs. Among these, H.264, widely recognized as AVC, has emerged as a cornerstone for video compression. Its exceptional ability to deliver high-quality video at relatively low bitrates has made it the codec of choice for streaming platforms, broadcast media, and video conferencing. By enabling efficient transmission of video content over networks with limited bandwidth, these algorithms ensure smooth and high-quality viewing experiences for users worldwide.

However, despite their efficiency, compression techniques often come with trade-offs, introducing visual artifacts such as blocking, mosquito noise, and posterization. These artifacts, while reducing file sizes, can degrade the overall quality of the video, potentially impacting user satisfaction. Addressing these challenges requires innovative solutions that not only optimize video delivery but also enhance the end-user experience.

This paper introduces cutting-edge techniques leveraging artificial intelligence (AI) to overcome the limitations of traditional video compression methods. Our approach combines state-of-the-art neural networks, trained within the Generative Adversarial Networks (GANs) framework, to enhance video quality in both archive materials and live streaming scenarios. As depicted in Fig. 2, GANs offer a powerful and flexible framework for generating realistic image details and restoring visual quality. By fine-tuning these networks, we achieve real-time performance on mid-range GPUs, making our solution accessible and scalable for a variety of applications. Furthermore, optimizations for mobile devices, leveraging hardware accelerators such as CoreML, the Neural Engine (iOS), and mobile GPU/WebGL support (Android), expand the reach of this technology to portable and web-based platforms.

Key Scientific Contributions

Our research introduces the following innovations:

1. **Advanced Loss Functions.** By combining perceptual metrics with signal-based evaluations, we achieve a balanced reconstruction of video frames that are both visually pleasing and technically accurate.
2. **Efficient Neural Network Architectures.** These architectures are designed to reduce computational overhead without sacrificing output quality, enabling real-time and faster-than-real-time processing.
3. **Novel GAN-Based Training Techniques.** These methods improve the generation of natural and realistic image details, minimizing artifacts and enhancing overall video quality.

Applications and Deployment

We have developed a suite of products that integrate these AI-powered enhancements directly into video players. Our solution processes video frames in real-time just before they are displayed, enabling the upscaling of video resolution and the removal of compression artifacts. This dual capability reduces the bandwidth required for streaming while restoring fine image details lost during compression. Moreover, the flexibility of these neural networks allows customization for specific types of content (like sports, cartoons etc.), or even individual titles. This adaptability ensures optimized video quality for diverse viewing experiences, achieving high-quality streaming at significantly reduced bitrates. The compact and efficient design of the networks also minimizes storage requirements, making

them ideal for deployment across a wide range of devices. Additionally, the technology enhances video conferencing experiences by improving visual clarity, even under challenging network conditions.

Conclusion

By integrating advanced AI techniques into the video processing pipeline, our approach delivers high-quality streaming experiences at reduced costs. With the ability to enhance video content in real-time, minimize compression artifacts, and adapt to a wide array of devices and use cases, this solution represents a significant advancement in video streaming technology. The improved efficiency and visual quality enabled by these innovations not only elevate user satisfaction but also position our technology as a valuable asset in the competitive video streaming industry.

1. Ledig C., Theis L., Huszár F., Caballero J., Cunningham A., et al. "Photo-Realistic Single Image Super-Resolution Using a Generative Adversarial Network." *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2017. DOI: 10.1109/CVPR.2017.19.

2. Shi W., Caballero J., Huszár F., Totz J., Aitken A. P., et al. "Real-Time Single Image and Video Super-Resolution Using an Efficient Sub-Pixel Convolutional Neural Network." *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016. DOI: 10.1109/CVPR.2016.207.

3. Chun D., Kim T. S., Lee K., Lee H. J. "Compressed Video Restoration Using a Generative Adversarial Network for Subjective Quality Enhancement." *IEIE Transactions on Smart Processing & Computing*, 2020, vol. 9, no. 1, pp. 1–6. DOI: 10.5573/IEIESPC.2020.9.1.001.

HATE SPEECH DETECTION: A HYBRID APPROACH

Houda EL Bouhissi¹, Tetiana Hovorushchenko², Olga Pavlova²

¹ *University of Bejaia, Faculty of Exact Sciences, LIMED Laboratory, Bejaia , 06000, Algeria*

² *Khmelnytskyi National University, Instytut's'ka str., 11, Khmelnytskyi, 29016, Ukraine*

With the explosion of social media, people are using more and more social media networks to communicate, buy products or post tweets. Social media, including Facebook and Twitter, are currently among the most widely used channels for communication. Certain social network posts may enclose hate speech, which can lead to "cyber conflict," which can have an impact on social life at the individual and countries' levels. Recently, in Algeria, the authorities to preserve the dignity and security of the citizens are exploring mechanisms to control the offensive messages posted by a minority or groups of people whose aim is to destabilize the country. In this paper, we provide a background on hate speech detection and its most related approaches. Furthermore, we present our work on detecting hate speech-language on Algerian tweets. Our proposal is based on domain ontologies.

Keywords: Discrimination, Hate Speech, Knowledge, Machine Learning, NLP, Ontology, Sentiment Analysis.

1 Introduction and Motivation

The growing presence of social media has brought up a multitude of possibilities and opportunities; we can for example discuss and communicate, buy or sell products, read the news and post tweets.

However, the large number of user's online interactions create new issues; and the content has become greatly difficult to manage.

The discrimination and public defamation of persons or groups have always been an important problem society has had to face. Furthermore, with the growth of social media, this issue lives more and more. Because people can express themselves freely and anonymously online and social media

platforms do not provide any oversight, hate speech is pervasive in online communication.

In this insight, it is essential to detect subjective information in the data set such as comments and tweets, to help for example businesses to understand the social sentiment of product, service, law, decisions while monitoring the online communications and reviews.

Recently, sentiment analysis of web content is becoming increasingly an important issue for developers due to augmented communication through Internet, such as social media, e-mail, and forums.

Regardless its effect, there is no common definition of hate speech; in general, organizations refer to it as the propagation of racism, sexism, or religious hate.

Many international locations have brought legal guidelines to oppose hate speech and discrimination that is a motive for “cyber conflict” and have an effect on social lifestyles both individual and country at the whole.

In Algeria, the authorities explore mechanisms and tools to detect hate speech in social media to punish the authors of the public defamation. The majority of internet sites also implement specific language restriction measures based on their own definitions.

Different approaches have been proposed to fight the dissemination of this kind of harmful content among the web due to the large volume of data generated continuously. This means, that users can report content directly to the platform provider, so that it can be evaluated and/or removed. In addition, bigger platforms outsource active content moderation to external companies.

The detection of hate speech is a very challenging problem for humans as well as for algorithms. In many cases even for humans, it is extremely difficult to distinguish hate speech from other content, which might be harmful, but does not fall under the category of hate speech.

However, current technologies are still reliably unable to identify hate speech. Therefore, new methods for automated hate speech identification and monitoring are required.

Recently, techniques for text mining and natural language processing (NLP) have considerably increased according to recent innovations in machine learning (ML).

Creative use of advanced artificial intelligence techniques, in particular, the use of ML has become increasingly popular to support content moderation in online platforms [1].

The aim of the present research is to analyse the current sentiment analysis-focused systems. Moreover, the development of an efficient knowledge-based system for hate speech discovery across several social media platforms, which allows the surveying of person's opinion on any matter that, is important to them or to the society.

The main contribution of this work is the detection of hate speech in Algerian tweets based on domain ontology. Furthermore, we propose a classification methodology for recognizing hate speech tweets.

The rest of the paper is organized as follows: Section 2 presents the background of the study, which covers overviews of the hate speech issue and domain ontologies. Section 3 presents the most important studies on the hate speech detection. Section 4 describes the proposed approach. Section 5 reports the results of the proposed methodology. Finally, Section 6 concludes the paper and outlines future projects.

2 Theoretical background

For a better understanding of the concepts, in this section, we present the theoretical background of the hate speech and our contribution. This includes brief overviews of the topic and the deep architectures used.

2.1 Hate Speech

Cyber-hate is the term used to describe hate speech and violent communication that takes place online [2]. It may be described as *"Any use of electronic communication technology to spread racist, religious, extremist, or terrorist messages"* and is a limited and particular type of cyberbullying.

Hate speech differs from cyberbullying in that it may have an impact on entire communities in addition to individuals [3].

In a general way, hate speech is a practice by a minority of persons or groups of persons, in verbal or textual form, involving words of discrimination or violence against other persons based on, for example, their race, nationality, religion, color, and sex.

Anis [4] argued that hate speech may take many forms, including insults, provocations, verbal abuse, and physical assault.

Chetty and Alathur [5] classified hate speech messages into the following categories:

- Gendered hate speech

This category encompasses any prejudice toward a certain gender or any devaluation of a person based on their gender. Misogyny in whatever form is also include.

Furthermore, Jha and Mamidi [6] stated that sexism can take one of two forms: hostile (which is an overtly hostile attitude), or benevolent.

- Religious hate speech

This includes any kind of religious discrimination, such as Islamic sects, calling for atheism, Anti-Christian and their respective denominations or anti-Hinduism and other religions.

Moreover, Albadi et al. [7] indicated that religious hate speech is considered as a motive of crimes in countries with highest social crimes.

- Racist hate speech

This category includes any racial offense sort or tribalism, regionalism, xenophobia (especially for migrant workers) and nativism (hostility against immigrants and refugees) and any prejudice against particular tribe or region.

For instance, offending an individual because he belongs to a particular tribe, region, country, or favoritism of a particular tribe. Add to that, offending the appearance and color of individual.

2.2 Ontology

The Semantic Web [8], which is an extension of the current Web, has as a major objective to bring semantics to the manipulated information in order to enable the communication between agents (humans and machines) with the use of metadata that describe the available information.

These metadata are provided by the critical component of the Semantic Web, the ontologies. An ontology is defined as "*an explicit, machine-readable specification of a shared conceptualization*" [9].

Ontologies represent a technology whose goal is to improve the organization, management and understanding of electronic information. They serve as a standardized vocabulary for knowledge discovery and sharing.

Ontologies have been used in a wide range of applications such as artificial intelligence and knowledge-based systems. Other applications involve Data integration, intelligent information integration, semantic-based access to the Internet and extracting information from texts.

Moreover, the most important contribution of ontologies is the key role they play in the development of the Semantic Web [10].

3 Literature Survey

On social media, harmful language is a complex issue with a diverse variety of overlapping forms and targets. A number of research has been conducted in the literature to identify abusive language, including hate speech, insulting language, and cyberbullying.

Recently, researchers have become more interested in identifying hate speech on social media. The majority of the studies that we found for hate speech were conducted for English. However, some other languages were considered.

Additionally, the literature has essentially modeled the issue as a supervised classification task. Here, we highlight the research on hate speech detection that is the most pertinent.

Kwok and Wang [11] proposed a new approach for detecting racist hate speech against black people in Twitter. The model uses Machine-Learning algorithms and can be improved with other features such as sentiments features.

Waseem and Hovey [12] presented a new methodology for detecting hate speech in Twitter platform. The proposed methodology focused on two types of hate speech: Racism and Sexism.

Djuric et al. [13] used word embedding to learn low-dimensional text embedding to classify users' comments as Harmful or clean.

Nobata et al. [14] proposed a new approach that incorporates different features such as linguistic and syntactic features to detect abusive language in user comments. The

results showed that combining all these features improve the performance of the detection process.

Davidson et al. [15] presented an empirical study, which explores different classifiers to categorize offensive and hate language in tweets. The authors constructed and published their own data set of about 24.802 tweets labelled into three categories: hate speech and offensive language and not.

Jaki and Smedt [16] conducted a study on over 50.000 right-wing German hate tweets posted during the German elections (between August 2017 and April 2018). The results of the analysis show how the insights can be employed for the development of automatic detection systems.

Bohra et al. [17] built a Hindi-English code-mixed dataset that involves online tweets posted on Twitter. The tweets are annotated and each tweet belongs to a different class (Hate Speech or Normal Speech). In addition, the authors propose a supervised algorithm for detecting hate speech in the text using different features such as characters and words.

Park and Fung [18] proposed a novel methodology, which combines two classifiers to categorize text. The first classifier categorizes the text to abusive language or not. Then the text labeled as abusive language is categorized into abusive language category such as sexist or racist.

Gambäck and Sikdar [19] introduce a deep learning-based Twitter hate-speech text classification system. Each tweet is categorized by the classifier within one of four predetermined categories : non-hate speech, racism, sexism, and both (racism and sexism). The authors employed four convolutional Neural Network models were based on semantic information. The model performed best, with higher precision than recall, and a 78.3% F-score.

Pitsilis et al. [20] built an ensemble of Long-Short Term Memory (LSTM)-based classifiers to detect hate speech content in Twitter. The experiment results showed that the proposed approach achieved better results.

Zhang et al. [21] proposed a new model that used neural network model to detect hate speech on Twitter. The authors constructed their own dataset by searching tweets related to religion and refugees in general. The results showed that the proposed approach achieved the highest F-score measure.

Badjatiya et al. [22] proposed a new methodology based on deep learning. To detect hate speech on text, the authors used three different models. The results of experiment showed that the combination of the LSTM model and the Gradient Boosting Decision Tree led to a best accuracy.

Modha et al. [23] presented a novel approach to detect and visualize online aggression on social media. The authors designed a user interface based on a web browser plugin over Facebook and Twitter to visualize the aggressive comments posted on the Social media user's timelines.

Salminen et al. [24] perform experiment with different classification algorithms on 197,566 comments collected from YouTube, Reddit, Wikipedia, and Twitter. The dataset involves 80% comments labeled as non-hateful and the remaining 20% labeled as hateful.

4 Proposed methodology

Hate speech is now a hot topic in the social media age. The Internet's anonymity and flexibility have made it simple for users to communicate aggressively.

Techniques that automatically detect hate speech are greatly needed as the volume of hate speech online rises. Additionally, the ML and NLP communities have been quite interested in these issues.

How to effectively identify the nature of a new posting (hateful or no)? To answer this question, our main goal is to propose a novel method that can improve the existing approaches and provide good performance/accuracy.

The purpose of this work is to detect hate speech in French social media regarding Algerian people. This section describes the proposed system, which we have employed to detect and classify tweets into different classes.

Our proposal is motivated by the lack of resources for Algerian corpora analysis and the interconnection of haltered and extremists contents

Fig. 1 illustrates the complete research methodology. As shown in this figure, the proposal involves five key steps namely:

(1) "Data Collection" which extracts useful text from posted tweets for building our dataset,

(2) "Data Preprocessing" includes a set of phases for tweets preprocessing, this step cleans the preprocessing tweets and keep only the cleaned tweets.

The third step (3) "Ontology matching" involves the ontology construction, which aims to classify tweets into hateful and not hateful,

(4) "Feature engineering" that to convert the tweet texts to numerical values for the use of ML algorithm that require numerical values.

Finally, the last step (5) "Tweets classification" in which the tweets are classified according to their category.

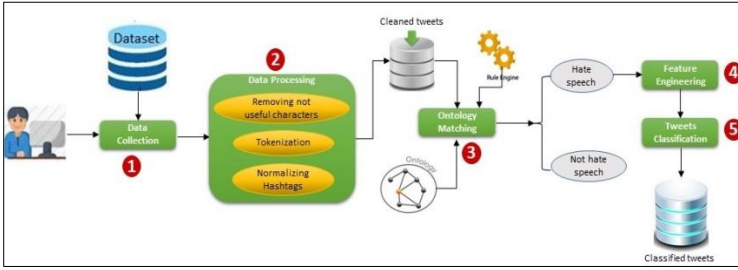


Fig. 1. System architecture.

We can describe our work in the whole as follows: starting from a set of posted tweets, written in French by several Algerian users, each tweet is first captured and preprocessed with NLP techniques. Then, a domain ontology is built to classify tweet into hateful or not according to hate speech terms collected from the Algerian vocabulary. The selected tweets are converted to numerical vector to apply a classification ML algorithm and finally, we classify the hateful tweets into three categories:

- Racist (Ra)
- Sexist (Se)
- and religious (Re).

Next, we introduce in detail the main steps of the proposed approach and describe how we use the extracted information to classify out tweets.

Step1: Data Collection

This step is very important for our methodology, hence, for our ML algorithm, we create a dataset from

Twitter regarding posts written in French. We extract then tweets from content about Algeria.

Note that, we added further words according to the Algerian tweets vocabulary (table1) and we labeled them according to a hate speech category.

Table 1. Examples of words used for tweets collection.

Hate speech category	Term/Hashtag	English meaning
Religion	Irhabi	Terrorist
Racism	Nigro	Black
Regional	Kabyle, Choui, Arbi	Kabyle, Chaoui, Arab
Obscene	Menteur	Liar
Violence	Les arabes sont soumis	Arabs are subordinate
Sexism	Comment une femme nous gouverne	How a woman governs us

Step 2: Data Preprocessing

In general, text preprocessing makes better classification results. For this purpose, we applied various preprocessing-techniques to clean noisy and non-informative features from the tweets.

Retrieved tweets mostly involved numerous words due to convention names, thus a preprocessing phase is needed. In this step, we clean the posted tweets in order to filter them and reduce workspace where we use NLP techniques to extract the useful words from the whole messages.

- **Removing not useful characters**

In this phase, we remove all useless characters that appears in the specific tweets, which includes punctuation marks and special characters (: , ; & ! ? n), numerical characters, stop words, emoji, punctuation and URLs and links that do not contribute towards any kind of sentiment in the text for the tweets.

In addition, words such as "isn't" and "I'll", are replaced by the corresponding expanded form. Finally, the result words are converted to lowercase.

- **Tokenization**

Information retrieved from the tweets contains different words. For simplifying our text, we use the tokenization process for splitting it into very simple items (called tokens) such as numbers, punctuation and words of various types. For this purpose, we use a set of rules to tokenize the words.

The tokenization converts each tweet into a set of tokens or items, and then the porter stemmer converts these tokens to their root forms, for example: classified to classify using porter stemmer.

For example:

"It is battery life is great. It is very responsive to touch. "

will be tokenized to:

"It" "is" "battery" "life" "is" "great" "It" "is" "very" "responsive"
"to" "touch".

Further, to reduce workspace, we delete similar tokens according to their denotation and specification.

- Normalizing Hashtags

Furthermore, we also normalized hashtags into words, for example, "#MuslimsAreTerrorist" became "Muslims are terrorist". This is because such hashtags are often used to compose sentences, and we require full words instead. For this purpose, we use a dictionary based look up to split such hashtags, which were made of coherent and correct usage of words to retrieve the context derived from these hashtags.

Step 3: Ontology matching

The most important step of the research presented here is to build an ontology to formalize online hate in textual Algerian content expressed in French. The built ontology highlights the main concepts for hate detection.

However, building an ontology from the scratch is a difficult task, therefore, using data extracted from sources with partial human expertise involved is an efficient way since the ontology constructed will meet the domain requirements.

From this viewpoint, a domain expert designed the ontology skeleton for storing existing Algerian corpora in a Knowledge Graph that is interoperable.

The Ontology construction approach follows the general methodology that involves three main phases [25]: "knowledge acquisition", "construction of the conceptual model" and finally the "consolidation".

The ontology will be populated by the useful information collected from hatebas.org¹, which is a collaborative, regionalized repository of multilingual hate speech.

After the ontology is created and populated, we perform a matching peer to peer between the ontology and the tweets to keep only the speech involving hateful content (we consider here only instances matching).

A suggested solution for the semantic heterogeneity issue is ontology matching. Finding correspondences between entities that are semantically associated is the goal. These correspondences can be leveraged to improve the provisional ontology and recognize if the tweets are harmful or not.

We employ a semantic similarity measure technique to distinguish between similar and dissimilar entities throughout the matching process [26]. When two entities are compared for semantic similarity, a numerical number (included in [0, 1]) indicating how similar they are to one another is returned.

Ontology matching commonly uses semantic similarity metrics. Based on the background knowledge in the ontologies annotations and massive text corpora, they are used to compare words, concepts, and ontology classes.

A simple similarity measure that is the most used metric is SimRada (see equ. 1), based on the shortest path between two nodes in the graph.

SimRada is equal to the inverse of the shortest path length between two concepts and can be defined as:

$$SimRada(x, y) = \frac{1}{distSP(x, y) + 1} \quad (1)$$

In addition, we define a threshold value (0.6), which represents the required accuracy. According to the matching result, we keep only the tweets with a matching degree higher than or equal to 0.6.

Therefore, this step is useful to filter the tweets and keep only the tweets that probably contain hate speech.

Step 4: Feature Engineering

This step is a preliminary step for the ML algorithm, which cannot understand the classification rules from the raw

¹ <https://hatebase.org/>

text. These algorithms need numerical features to understand classification rules. Hence, in text-classification one of the key steps is feature engineering.

This step is used for extracting the key features from raw text and representing the extracted features in numerical form. In this study, we have performed the Doc2vec feature engineering technique to produce a numeric vector representation of each tweet.

Step 5: Tweets Classification

Teaching an automated system to recognize hate speech is a challenging task. ML methods are now applied widely across life sciences to develop predictive models.

The goal of this step is to classify the hate speech tweets in their appropriate category.

In this step, we use Support Vector Machine (SVM) algorithm, which is the most popular algorithm used to perform classification [27]. Usually, the classifier goes through two main phases, which are training and testing.

The classifier will be fed training data throughout the training phase in order to build a model. The testing phase will next employ this model to forecast the labels of test data.

The result of this step is a classification of the harmful tweets into three categories:

- Racist (Ra)
- Sexist (Se)
- Religious (Re).

5 Experiment and evaluation

To validate the proposed approach and gain insights about its fullness and perfection, we implemented a software tool in python.

In this section, we presented our ML model, which is trained and tested on the respective dataset. We performed two experiments using the SVM algorithm: the first using domain ontology and the second without domain ontology.

We analyze about 9500 tweets, the results of the proposed classification is described as follow in Table 2 :

Table 2. Tweets classification.

Hate speech category	Tweets (%)
Category 1	23.59
Category 2	43.29
Category 2	33.12

To evaluate the efficiency of our approach, we use metrics that are widely used: "Precision" and "Recall" [29] [30]. Both precision and recall are based on an understanding and measure of relevance.

- Precision also called positive predictive value (see equ.2) is the ratio of relevant classified tweets among the classified tweets :

$$Precision = \frac{\text{Correctly classified tweets}}{\text{Total classified tweets}} \quad (2)$$

- Recall also known as sensitivity (see equ. 3) is the ratio of relevant classified tweets that have been retrieved over the total amount of classified tweets :

$$Recall = \frac{\text{Correctly classified tweets}}{\text{Total useful classified tweets}} \quad (3)$$

Table 3 describe the precision and recall of each experiment performed.

Table 3. Evaluation Metrics

Metric	Result With ontology (%)	Result Without ontology (%)
Recall	75.00	55.00
Precision	83.00	60.00

Table 3 shows that the combination of the SVM algorithm and the ontology provide best results.

6 Conclusion

Automated systems for detecting hate speech have received a lot of attention in recent years. Facebook and Google attempted to detect hate speech using artificial intelligence.

However, there is still a lot of work that needs work by humans since the systems in place are insufficiently strong.

In this paper, we provided an overview of the most important works that deal with hate speech and presented a novel knowledge-based methodology for identifying and tracking hate speech in Algerian Tweets based on domain ontologies. The use of domain ontologies improve significantly the efficiency of our system. Our proposal is supported by a software-tool, and preliminary results demonstrate its effectiveness.

For future work, we plan to enhance the precision by applying different classifications algorithms such as Logic Regression, and to expand the test set to include Arabic and Amazigh languages (the second official language in Algeria) tweets from multiple regions to cover different dialects and cultures.

1. Bekka, R., Kherbouche, S., & El Bouhissi, H. Distraction Detection to Predict Vehicle Crashes: A Deep Learning Approach. *Computación y Sistemas*, 26(1) (2022).
2. Miro-Linares, F. and Rodriguez-Sala, J. J. "Cyber hate speech on twitter: Analyzing disruptive events from social media to build a violent communication and hate speech taxonomy," *Int. J. Des. Nat. Ecodynamics*, vol. 11, no. 3, pp. 406–415 (2016).
3. Brown, A. "What is hate speech? Part 1: The Myth of Hate," *Law Philos.*, vol. 36, no. 4, pp. 419–468 (2017).
4. Anis, M. Y. and Maret, U. S. "Hatespeech in Arabic Language," in *International Conference on Media Studies* (2017).
5. Chetty, N. and Alathur, S. "Hate speech review in the context of online social networks," *Aggress. Violent Behav.*, vol. 40, no. May, pp. 108–118 (2018).
6. Jha, A. and Mamidi, R. "When does a compliment become sexist? Analysis and classification of ambivalent sexism using twitter data," *Proc. Second Work. NLP Comput. Soc. Sci.*, pp. 7–16 (2017).
7. Albadi, N., Kurdi, M. and Mishra, S. "Are they Our Brothers? Analysis and Detection of Religious Hate Speech in the Arabic Twittersphere," 2018 *IEEE/ACM Int. Conf. Adv. Soc. Networks Anal. Min.*, pp. 69–76 (2018).
8. Berners-Lee, T., Hendler, J. and Lassila, O. *The Semantic Web. Scientific American*, 284(5) :34 – 43 (2001).
9. Gruber, T. A translation approach to portable ontology specifications, *Knowledge Acquisition* 5(2), pages 199-220 (1993).

10. El Bouhissi, H., Salem, A-B.M. and Tari, A. ‘Semantic enrichment of web services using linked open data’, *Int. J. Web Engineering and Technology*, Vol. 14, No. 4, pp.383–416 (2019).
11. Kwok, I.;Wang, Y. Locate the hate: Detecting tweets against blacks. In *Proceedings of the Twenty-seventh AAAI Conference on Artificial Intelligence*, Washington, DC, USA, 14–18 (2013).
12. Waseem, Z.; Hovy, D. Hateful Symbols or Hateful People? Predictive Features for Hate Speech Detection on Twitter. In *Proceedings of the NAACL Student Research Workshop*, San Diego, CA, USA pp. 88–93 (2016).
13. Djuric, N.; Zhou, J.; Morris, R.; Grbovic, M.; Radosavljevic, V.; Bhamidipati, N. Hate Speech Detection with Comment Embeddings. In *Proceedings of the 24th International Conference on World Wide Web—WWW ’15 Companion*, Florence Italy pp. 29–30 (2015).
14. Nobata, C.; Tetreault, J.; Thomas, A.; Mehdad, Y.; Chang, Y. Abusive Language Detection in Online User Content. In *Proceedings of the 25th International Conference on World Wide Web—WWW ’16*, Montreal, QC, Canada pp. 145–153 (2016).
15. Davidson, T.; Warmesley, D.; Macy, M.; Weber, I. Automated Hate Speech Detection and the Problem of Offensive Language. *arXiv* (2017).
16. Jaki, S., & Smedt, T.D. Right-wing German Hate Speech on Twitter: Analysis and Automatic Detection. *ArXiv*, abs/1910.07518 (2019).
17. Bohra, A. & Vijay, D. & Singh, V. & Akhtar, S.S. & Shrivastava, M. A Dataset of Hindi-English Code-Mixed Social Media Text for Hate Speech Detection. 36-41. 10.18653/v1/W18-1105 (2018).
18. Park, J.H.; Fung, P. One-step and Two-step Classification for Abusive Language Detection on Twitter. In *Proceedings of the First Workshop on Abusive Language Online*, Vancouver, BC, Canada pp. 41–45 (2017).
19. Gambäck, B.; Sikdar, U.K. Using Convolutional Neural Networks to Classify Hate-Speech. In *Proceedings of the First Workshop on Abusive Language Online*, Vancouver, BC, Canada pp. 85–90 (2017).
20. Pitsilis, G.K.; Ramampiaro, H.; Langseth, H. Effective hate-speech detection in Twitter data using recurrent neural networks. *Appl. Intell.*, 48, 4730–4742 (2018).
21. Zhang, Z.; Robinson, D.; Tepper, J. Detecting Hate Speech on Twitter Using a Convolution-GRU Based Deep Neural Network.

22. Badjatiya, P.; Gupta, S.; Gupta, M.; Varma,V. Deep Learning for Hate Speech Detection in Tweets. In Proceedings of the 26th International Conference on World Wide Web Companion, Perth, Australia pp. 759–760 (2017).
23. Modha, S., Majumder, P., Mandl, T. and Mandalia, C. Detecting and Visualizing Hate Speech in Social Media: A Cyber Watchdog for Surveillance. Expert Systems with Applications. 161. 113725. 10.1016/j.eswa.2020.113725 (2020).
24. Salminen, J. & Hopf, M. & Chowdhury, S. & Jung, S. & Almerakhi, H. & Jansen, J. Developing an online hate classifier for multiple social media platforms. 10. 1. 10.1186/s13673-019-0205-6 (2020).
25. EL Bouhissi, H. and Malki, M. Semantic Web Service Ontology Construction: A Reverse Engineering Approach. In the 11th ACS/IEEE International Conference on Computer Systems and Applications (AICCSA'2014), November 10-13, Doha, Qatar (2014).
26. EL Bouhissi, H. , Malki , M. and Sidi Ali Cherif, M.A. From User's Goal to Semantic Web Services Discovery: Approach Based on Traceability. In International Journal of Information Technology and Web Engineering (IJITWE), 9(3), 15-39, July-September (2014).
27. Dakhaz, A. Machine Learning Applications based on SVM Classification: A Review (2021).
28. Mustafa Abdullah, D. and Mohsin Abdulazeez, A. Machine Learning Applications based on SVM Classification A Review (2021).
29. El Bouhissi, H., Adel, M., Ketam, A., & Salem, A. B. M. Towards an Efficient Knowledge-based Recommendation System. In IntellTSIS (pp. 38-49) (2021).
30. Sokolova, M. and Lapalme, G. “A systematic analysis of performance measures for classification tasks,” Inf. Process. Manag., vol. 45, no. 4, pp. 427–437 (2009).

FINDING DEFECTS IN PERIODIC SCTRUCTURES

Volodymyr Zhydkov, ORCID: 0000-0003-0553-6384

*V.M. Glushkov Institute of Cybernetics of the NAS of Ukraine, Kyiv,
Ukraine*

lynx233@yahoo.com

Rejection of defective parts and elements is an important part in quality improvement, maintenance and accident prevention. Some imaging instrumental methods of test and inspection have improved ability to reveal defects, but currently have disadvantages, such as low accuracy of defect classification

and need of human operator. The paper presents generally applicable algorithmic solution for a system that can analyze images and reveal defects automatically on par or better than a human operator. The general principle is based on approach that some general properties about the analyzed part (periodicity) is known beforehand; thus, image of the object examined in its defect-less state can be reconstructed and, by comparing image to the reconstruction, the defects can be revealed.

Keywords: regular structures, defectoscopy, image processing.

Introduction

One of the primary issues with nearly all modern methods of defect detection by images is that they should all have a catalog of known defect types, upon which algorithms supposed to be adjusted/trained. However, there are inherent disadvantages to such approach, such as possible appearance of previously unknown defect types. Therefore, there was tested another approach: attempt to recognize/match not the defects, but the background upon which defects are superimposed.

Statistical Model of the Problem

Let matrix $M \in \mathbb{R}^{n \times m}$ be some experimentally obtained two-dimensional image of size $n \times m$, corresponding to the structure studied, and matrix $L \in \mathbb{R}^{n \times m}$ be a two-dimensional image of size $n \times m$ as well, which corresponds to an idealized model of the structure that does not contain defects. The processing is based on assumption that image M obtained experimentally, it does not contain systematic measurement errors, and the noise in the image M is only white. It follows from this assumption that $D = M - L$ is the deviation of the image of a correctly constructed idealized model from the measured image should have a normal distribution density of element values corresponding to white noise, with an accuracy of up to the p , that is defects to image area ratio. This provides a criterion for determining whether a certain constructed image L corresponds to experimental data M . Following the ab initio principle of maximum likelihood, we consider the weight function

$$W(M - L) = \sum_{i=1}^n \sum_{j=1}^m \ln \left(1 - \left(1 - e^{-(M(i,j) - L(i,j))^2 / 2\sigma^2} \right) \times (1 - p) \right), \quad (1)$$

which determines how closely the image L corresponds to the image M . We also assume that the noise of the image D has a normal distribution with standard deviation σ . For practical purposes, this weighting function can be replaced by the "truncated" L_2 -norm, which is calculated according to the expression

$$W_c(D) = \sum_{i=1}^n \sum_{j=1}^m \min(D^2(i, j), -2\sigma^2 \ln(p)), \quad (2)$$

and is quite close to (1).

To take the above into account when filling an image with translation symmetry in one direction, we will use the formula

$$L = Q_1 + F(Q_2) \times Q_3, \quad (3)$$

where $\times$ is element-wise matrix multiplication operation, Q_l , $l = 1, 2, 3$, are matrices of the same size as the original image, filled according to the formula

$$Q_l(i, j) = a_i + b_i i + c_j + d_i j + e_i i^2 + f_j j^2 + g_i i^2 j + h_i i j^2 + k_i i^2 j^2. \quad (4)$$

Similarly, for a model with symmetry in two directions, one can use

$$L = Q_1 + F_1(Q_2) \times F_2(Q_3) \times Q_4, \quad (5)$$

where Q_l , $l = 1, 2, 3, 4$, are matrices filled according to formula (4), and F_1 and F_2 are periodic functions applied to the elements of the matrix.

For an image with one direction of symmetry, we can formulate the following optimization problem:

$$\Phi^* = \Phi(a_1^*, \dots, k_3^*) = \min_{(a_1, \dots, k_3) \in \mathbb{R}^{27}} W(M - Q_1 + F(Q_2) \times Q_3). \quad (6)$$

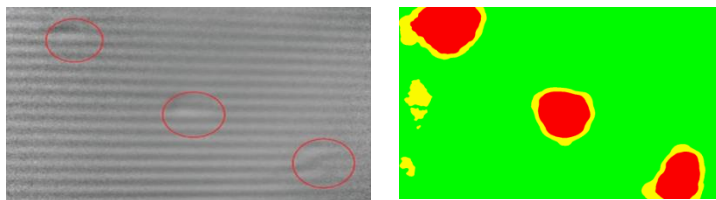
where coefficients $a_l, \dots, k_l$ define l -th ratio (4), W is norm used, F is used periodic function.

Similarly, for two directions of symmetry, we can formulate an optimization problem:

$$\Phi^* = \Phi(a_1^*, \dots, k_4^*) = \min_{(a_1, \dots, k_4) \in \mathbb{R}^{36}} W(M - Q_1 + F_1(Q_2) \times F_2(Q_3) \times Q_4), \quad (7)$$

where F_1, F_2 are approximated periodic functions.

The method was implemented in Octave language and tested on real-life images. The source of the image in the case is shearography of object under heat flux. Such images are rather informative, but difficult to analyze, and the algorithm can help with that.



On the left – input test image, visible defects are outlined with red ellipses, on the right – algorithm output, highlighting defective areas. Here, red corresponds to defect probability over 99%, yellow to defect probability 90-99%. As one can see here, the method allows to detect defects not really visible by naked eye and can work reliably enough.

Conclusions

The method of revealing defects based on background recognition is presented. It's applicability to real life data demonstrated. It is shown, that it can reliably reveal defects having minimal information about object studied, with no library of defects, previous training or human intervention. It can be implemented and be useful in automated defectoscopy solutions.

1. Zhydkov V.A. Revealing defects in periodic structures. Computer Mathematics. 2017. No 2. P. 21–29.
2. Stetsyuk P.I., Savitsky V.V. On Defects Searching in regular 3D-structures. Journal of Automation and Information Sciences. 2018. Vol. 50. Issue 3. P. 21–37

PROTECTION OF MULTILAYER NETWORK SYSTEMS FROM TARGETED GROUP ATTACKS ON THE PROCESS OF INTERSYSTEM INTERACTIONS

Olexandr Polishchuk ⁰⁰⁰⁰⁻⁰⁰⁰²⁻⁰⁰⁵⁴⁻⁷¹⁵⁹,
Dmytro Polishchuk ⁰⁰⁰⁰⁻⁰⁰⁰²⁻⁷⁸⁵²⁻⁰³⁷⁸,

*Pidstryhach Institute for Applied Problems of Mechanics and
 Mathematics, National Academy of Sciences of Ukraine, Lviv,
 Ukraine
 od_polishchuk@ukr.net*

Effective scenarios for targeted group attacks on the structure and processes of intersystem interactions are proposed. These scenarios are based on the concepts of aggregate-networks of

multilayer network systems (MLNS), which describe such interactions, and on the application of structural and flow cores of these aggregate-networks.

Keywords – multilayer network system, flow, aggregate-network, core, targeted attack

Introduction

Every real-world system is open, meaning it interacts with numerous other systems. In the theory of complex networks (TCN), the structure of such interactions is described by multilayer networks (MLNs) of various types [1]. Each layer of MLN represents the structure of separate system, while interlayer connections reflect the structure of intersystem interactions. One of the important problem studied within TCN is the development of strategies to protect MLNS from targeted attacks of various types. The complexity of this problem lies in the fact that the disruption of even one layer of such formation negatively impacts all layers connected to it. In TCN, scenarios for sequential attacks on the most important MLN nodes, identified based on specific criteria, are primarily developed, as their disruption would cause the greatest damage to the multilayer network [2]. The usefulness of such scenarios lies in their ability to model the most probable development of attack, thereby identifying the system components requiring prioritized protection. At the same time, group attacks on multilayer systems are significantly more dangerous, both in terms of ensuring protection and mitigating the consequences of disruptions. It should also be noted that destabilizing or halting the MLNS operation is possible even without damaging its structure. This study proposes effective scenarios for sequential group attacks on the structure and operation process of multilayer network systems.

Targeted group attacks on MLNS structure

Multilayer network is fully described by an adjacency matrix $\mathbf{A}^M$, where M is the quantity of MLN layers. The diagonal blocks of this matrix represent the structure of intrasystem interactions, while the off-diagonal blocks describe intersystem interactions. The elements of $\mathbf{A}^M$ are equal to 1 if a connection exists between the corresponding nodes of MLN and 0 otherwise. The **aggregate-network** of MLN is defined as the projection of all its layers onto one network layer. The adjacency matrix of this weighted network can be easily derived from $\mathbf{A}^M$ [3]. To identify groups of the most important MLN nodes, we use

the concept of k^{ag} -core of the aggregate-network. This is a subnetwork whose nodes have a generalized degree (the sum of input and output intra- and intersystem connections) at least $k^{ag}>1$ [4]. Clearly, the higher the value k^{ag} , the more structurally significant the group of MLN nodes identified by the k^{ag} -core becomes. Thus, scenario 1 of sequential targeted group attack on the MLN structure can be described as follows:

- 1) set the value $q = \max \{k^{ag}\}$;
- 2) create a list of groups of nodes that are part of q -core in MLN aggregate-network;
- 3) sort the compiled list of groups in descending order according to the selected importance indicator in aggregate-network, for example, the generalized structural degree of the group;
- 4) remove the first group from sorted list;
- 5) if the attack success criterion, for example, dividing the aggregate-network into unconnected components, is met, terminate the execution of scenario; otherwise, proceed to point 6;
- 6) if the list of groups with current value q is not exhausted, return to point 3; otherwise, proceed to point 7;
- 7) set $q=q-1$; if q is less than the minimum k^{ag} value, terminate the execution of scenario; otherwise, return to point 2.

Targeted group attacks on MLNS operation process

The operation process of multilayer network system is described using a flow adjacency matrix $\mathbf{V}^M(t)$, $t>T$ [3]. The structure of this matrix is identical to structure of the adjacency matrix $\mathbf{A}^M$. The elements of $\mathbf{V}^M(t)$ represent the normalized (to 1) flow volumes that have passed through the corresponding MLNS edges over a time interval of duration T . The flow aggregate-network of multilayer system has a structure identical to structure of the MLN aggregate-network. The elements of adjacency matrix of the flow aggregate-network are equal to the total flow volumes that have passed through the corresponding edges across all layers of multilayer system, also normalized to 1. To identify the most functionally significant groups of MLNS nodes, we introduce the concept of **flow λ_{ag} -core** of the multilayer system's aggregate-network. This core represents a subsystem in which the elements of flow adjacency matrix have values no less than $\lambda_{ag} \in [0, 1]$. Clearly, the larger the value λ_{ag} , the more functionally important group of MLNS nodes identified by

the λ_{ag} -core becomes. Thus, scenario 2 of sequential targeted group attack on the process of intersystem interactions can be described as follows:

- 1) set the value $\lambda_{ag} = 1$;
- 2) create a list of groups of nodes that are part of λ_{ag} -core in the MLNS flow aggregate-network;
- 3) sort the compiled list of groups in decreasing order based on the strength of their interaction [3] with MLNS aggregate-network;
- 4) remove the first group from sorted list;
- 5) if the attack success criterion, for example, reducing flow volumes in MLNS by a given amount, is met, terminate the execution of scenario; otherwise, proceed to step 6;
- 6) if the list of groups with the current value λ_{ag} is not exhausted, return to step 3; otherwise, proceed to step 7;
- 7) set $\lambda_{ag} = \lambda_{ag} - \delta$, where $0 < \delta < 1$ is a predefined value, for example, $\delta = 0,1$; if λ_{ag} is less than its minimum value for the flow aggregate-network, terminate the execution of scenario; otherwise, return to step 2.

Using the example of real-world network system, specifically the railway transportation system of the western region of Ukraine, it is demonstrated that scenario 2, based on the use of flow core of aggregate-network, is three times more effective in terms of the dimension of targeted group compared to scenario 1.

Conclusions

Based on structural and flow models of multilayer network system, methods for selecting the most important groups of nodes within the MLNS aggregate-network, both in terms of structure and functioning, have been proposed. Using the determined importance indicators of aggregate-network cores, effective scenarios for targeted group attacks on multilayer systems were developed. It has been shown that the flow-based approach to constructing such scenarios is significantly more effective than the structural ones.

[1] G. Bianconi, Multilayer Networks: Structure and Function, Oxford Univ. Press, Oxford, 2018.

[2] R. Kumar, A. Singh, Robustness in multilayer networks under strategical and random attacks, *Procedia Computer Science*, 173 (2020) 94-103.

[3] O. Polishchuk, Protection of Multilayer Network Systems from Successive Attacks on the Process of Intersystem Interactions, *CEUR-WS*, 3790 (2024) 545-557.

[4] Y.X. Kong et al, k-core: Theories and applications, *Physics Reports*, 832 (2019) 1-32.

ВИКОРИСТАННЯ СЕМАНТИЧНИХ МЕРЕЖ ДЛЯ ВИЯВЛЕННЯ ПОТЕНЦІЙНИХ ФЕЙКІВ

А.О. Болдак¹ [0000-0001-7151-0743] (boldak@wdc.org.ua),

К.В. Єфремов¹ [0000-0003-3495-6417]

О.В. Стусь¹ [0000-0003-3426-5093]

О.О. Дмитренко^{2,1} [0000-0001-8501-5313]

¹ *Національний технічний університет України "Київський
політехнічний інститут імені Ігоря Сікорського"*

² *Інститут проблем реєстрації інформації НАН України*

*boldak@wdc.org.ua, k.yefremov@wdc.org.ua,
o.stus@kpi.ua, dmytrenko.o@gmail.com*

Робота присвячена методам формування семантичних мереж для новинних повідомлень з метою виявлення потенційних фейків. Основна ідея полягає у використанні методики побудови семантичних мереж на основі ключових термінів. Розглядаються техніки обробки тексту, виділення ключових термінів і встановлення взаємозв'язків між ними. Описується метрика для вимірювання семантичної близькості повідомлень за допомогою міри Фробеніуса, що дозволяє оцінювати схожість семантичних мереж, що відповідають цим повідомленням. Це підвищує точність аналізу контенту та допомагає визначати надійні інформаційні джерела для подальшої валідації фактів. Використання цих методів сприяє підвищенню ефективності аналізу інформації, а також допомагає виявляти інформаційні впливи новинному просторі.

Ключові слова: обробка природної мови, семантична мережа, семантичний аналіз, виявлення фейкових новин, інформаційна безпека.

Вступ

В сучасному інформаційному просторі проблема дезінформації набуває все більшої актуальності, особливо з огляду на швидкість розповсюдження новин у медіа-потоків. Медіа-ресурси стають не лише каналами для передачі інформації, але й потенційними інструментами впливу, зокрема через поширення неправдивих або маніпулятивних відомостей. Це ставить перед науковцями завдання розробки ефективних методів виявлення та аналізу таких явищ.

Для виявлення неправдивих повідомлень широко використовуються методи обробки природної мови (NLP), зокрема ті, що базуються на машинному навчанні [1]. Ці підходи передбачають аналіз текстових характеристик новин з використанням алгоритмів класифікації, таких як нейронні мережі або дерева рішень. Однак, одним з основних недоліків таких методів є необхідність регулярного перенавчання моделей для адаптації до нових тем, змін у лексиці або появи нових мовних патернів. Це може уповільнювати процес обробки інформації та знижувати ефективність систем в умовах швидкого оновлення медіа-ресурсів [2, 3].

Одним із підходів, що розглянутий у цій роботі, є використання семантичних мереж для аналізу текстових повідомлень. Сучасні дослідження в цій сфері фокусуються на створенні моделей, що базуються на ключових термінах та іменованих сутностях, які дозволяють структурувати інформаційний потік та, зокрема, виявляти приховані смислові зв'язки між текстами. Важливу роль відіграють методи попередньої обробки тексту, виділення ключових термінів та аналізу взаємозв'язків між ними.

У цьому контексті, дане дослідження пропонує новий підхід до формування семантичних мереж, який дозволяє не лише структурувати текстовий контент, але й ефективно виявляти потенційно дезінформаційні джерела через аналіз їхньої семантичної близькості до інших новинних повідомлень.

Метою цієї роботи є розробити підхід до ідентифікації потенційних фейкових новин на основі порівняння семантичних мереж інформаційних повідомлень із сформованим шаблоном, що базується на достовірних джерелах.

Підхід до ідентифікації фейкових новин

У цій роботі запропоновано підхід до ідентифікації фейкових новин, заснований на використанні інформації з надійних і перевірених джерел, таких як провідні інформаційні агентства. Основна концепція полягає в тому, що фейкові новини визначаються на основі відмінностей між ними та достовірними повідомленнями. Для цього спочатку відбирається група новин, які за змістом є схожими на тестове повідомлення, але походять із достовірних джерел. На основі цієї вибірки формується шаблон – семантична мережа, яка служить еталоном для подальшого порівняння зі семантичною мережею тестового повідомлення. Для побудови семантичної мережі застосовується методика побудови семантичної мережі на основі PoSTagging, що представлена у наступних роботах [4, 5].

Припускається, що новини фейкового характеру, порівняно з достовірним контентом, мають суттєві розбіжності зі сформованим еталоном. Відповідно до цього припущення, розроблено метод порівняння новин з попередньо створеним еталоном патерном, який базується на даних з надійних джерел. Аналіз подібностей і відмінностей між тестовими повідомленням та еталоном дозволяє ідентифікувати новини, що можуть бути фейковими. Запропонований підхід дозволяє ідентифікувати недостовірні повідомлення серед великого потоку інформації та позначати їх для подальшої перевірки.

Додатково, цей підхід дозволяє визначати достовірність джерел інформаційних повідомлень. Порівняння тестових новин із заздалегідь сформованими еталоновими повідомленнями сприяє не лише виявленню фейкових новин, але й виявленню джерел, що систематично поширюють недостовірну інформацію. Це дозволяє створити систему, здатну фільтрувати інформаційний потік та підвищувати загальну якість новинного контенту.

Висновки

У цій роботі запропоновано підхід для ідентифікації потенційних фейкових повідомлень шляхом порівняння текстового контенту з достовірними джерелами інформації. Основною перевагою цього методу є відсутність необхідності регулярного перенавчання моделей, що робить його стабільнішим і менш залежним від змін тематики новин або лексики. Завдяки тому, що аналіз здійснюється на основі

змістовних відмінностей, цей підхід менше піддається впливу стилістичних варіацій у фейкових повідомленнях, що дозволяє фокусуватись на ключових фактах.

Хоча ефективність такого підходу не перевищує методи на базі машинного навчання, його ключова перевага полягає в швидкості ідентифікації, що критично важливо для обробки новин у реальному часі в динамічних інформаційних системах.

Також цей підхід дозволяє оцінювати надійність джерел, що поширюють інформацію, допомагаючи виявляти систематичних постачальників фейкових новин. Надійні джерела можуть бути інтегровані в базу фактів для покращення ефективності боротьби з дезінформацією та підвищення загальної якості новинного контенту.

1. Ahmed, A. A. A., Aljabouh, A., Donepudi, P. K., & Choi, M. S. (2021). Detecting fake news using machine learning: A systematic literature review. arXiv preprint arXiv:2102.04458.

2. Endsley, M. R. (2018). Combating information attacks in the age of the internet: New challenges for cognitive engineering. *Human Factors: The Journal of Human Factors and Ergonomics Society*, 60(8), 1081–1094.

3. Roozenbeek, J., & van der Linden, S. (2019). The fake news game: Actively inoculating against the risk of misinformation. *Journal of Risk Research*, 22(5), 570–580.

4. Lande, D. V., & Dmytrenko, O. O. (2020). Methodology for Extracting of Key Words and Phrases and Building Directed Weighted Networks of Terms with Using Part-of-speech Tagging. Selected Papers of the XX International Scientific and Practical Conference "Information Technologies and Security" (ITS 2020) In CEUR Workshop Proceedings (ceur-ws.org). - Vol-2859 (pp. 168-177). ISSN 1613-0073

5. Lande D.V., & Dmytrenko O.O. (2021). Using Part-of-Speech Tagging for Building Networks of Terms in Legal Sphere. In Proceedings of the 5th International Conference on Computational Linguistics and Intelligent Systems (COLINS 2021). Volume I: Main Conference. Kharkiv, Ukraine, April 22-23, 2021. CEUR Workshop Proceedings (ceur-ws.org). - Vol-2870 (pp. 87-97). ISSN 1613-0073.

ПОБУДОВА ПРОФІЛЯ АНДРАГОГА НА ОСНОВІ ВІДКРИТИХ ДЖЕРЕЛ ІНФОРМАЦІЇ

Юлія Рогушина¹, ORCID: 0000-0001-7958-2557,

Анатолій Гладун², ORCID: 0000-0002-4133-8169,

Сергій Прийма³, ORCID: 0000-0002-2654-5610,

Олена Аніщенко⁴, ORCID: 0000-0002-6145-2321.

¹ *Інститут програмних систем НАН України, Інститут
цифровізації освіти НАПН*

² *Міжнародний науково-навчальний центр інформаційних
технологій та систем НАН та МОН України*

³ *Таврійський державний агротехнологічний університет імені
Дмитра Моторного*

⁴ *Інститут педагогічної освіти і освіти дорослих
імені Івана Зязюна НАПН України
ladamandraka2010@gmail.com, glanat@yahoo.com,
pryima.serhii@tsatu.edu.ua, anishchenko.olena@gmail.com*

Розглянуто мету та засоби побудови профілю андрагога в системах інформаційної підтримки освіти дорослих. Визначено елементи його структури, методи семантичної обробки та зовнішні відкриті джерела.

Ключові слова: профіль андрагога, відкриті джерела інформації, семантичні вікі.

Вступ

Для програмної реалізації інформаційної підтримки діяльності андрагога потрібно побудувати формальну модель, що відображає структуру та взаємні відношення всіх основних суб'єктів та об'єктів екосистеми навчання дорослих та забезпечує актуальність представлених відомостей. Для деяких елементів цієї екосистеми вже створено різноманітні стандарти та схеми опису метаданих (наприклад, для опису навчальних об'єктів (Learning Objects, LO) та для профілювання здобувачів освіти), які враховують специфіку андрагогіки або можуть бути розширені відповідно до потреб цієї предметної області [1]. Але такий важливий суб'єкт цієї системи, як андрагог, залишається поза фокусом уваги в цих дослідженнях. У застосовних системах, що надають андрагогічні сервіси, він розглядається лише як суб'єкт, тобто як користувач зі специфічними повноваженнями, при цьому властивості самого андрагога практично не

фіксуються. Це значно обмежує можливості персоніфікації освітнього процесу (може враховуватися тільки історія взаємодії цього користувача з системою). Крім того, інші користувачі не можуть виконувати пошукові запити, в яких вказуються характеристики андрагога, саме через те, що ці характеристики не структуровані та не формалізовані. Важливо також, що представлена в системі інформація може бути застарілою або не об'єктивною.

Профіль андрагога: призначення та структура

Існуючі програмні системи, що використовуються для підтримки освітнього процесу, зазвичай враховуються лише базові дані про освіту і кваліфікацію викладача, без акценту на його спеціалізацію та практичний досвід у конкретній галузі. Але дорослі здобувачі мають більш глибокі та структуровані вимоги до рівня знань і досвіду андрагога відповідно до своїх уявлень про подальше використання результатів навчання. Важливо також поповнювати ці відомості на основі відкритих джерел, таких як наукометричні бази даних.

Приклади запитів, в яких такі додаткові характеристики могли б використовуватися: пошук андрагога, що має досвід викладання певного навчального курсу більше N років в освітніх установах певної країни; пошук LO з навчальних курсів А, що підготовлені андрагогом, який має досвід викладання навчального курсу В; відбір групи андрагогів, які можуть викладати дисциплін А, В і С, мають додаткові компетенції в сфері Х (наприклад, знають певну іноземну мову) та психологічно сумісні між собою; вибір групи незалежних експертів, що публікували за останні 5 років наукові дослідження в області В та не мають спільних публікацій.

Ці запити можуть бути використані для досягнення таких цілей, як формування команд інструкторів (з визначеною задачею та вимогами до учасників), створення груп викладачів та студентів для виконання грантів та дослідницьких проектів з урахуванням умов їх здійснення (наприклад, у військовій сфері), пошук незалежних експертів для оцінювання результатів навчання (з дотриманням умов неперетину інтересів), вибір консультантів для здобувачів освіти у складних предметних областях (наприклад, наукових керівників) тощо.

Для побудови структури такого профілю виникає потреба у співпраці експертів з освіти дорослих, які можуть визначити

найбільш значущі характеристики та їх можливі властивості, та спеціалістів з аналізу знань, яким потрібно знайти пертинентні моделі та семантичні технології для їх формалізації. В процесі побудови профілю андрагога потрібно розробити методи перетворення природномовних описів на чіткі визначення, що можуть бути однозначно інтерпретовані та придатні для автоматичної обробки.

Основна вимога до створення такого профілю полягає в тому, щоб зберігати базові, неопрацьовані значення характеристик, а всі похідні від них оцінки обчислювати відповідно до умов конкретної задачі. Наприклад, доцільно не обчислювати узагальнений рейтинг андрагога, а визначати його для конкретного курсу або дисципліни. Важливими складовими профілю андрагога мають бути відношення, які формалізують семантику його зв'язків з іншими об'єктами та суб'єктами екосистеми освіти дорослих. Наприклад, андрагог може бути пов'язаний з іншими андрагогами відношеннями “співавтор” або “співробітник”, зі здобувачами освіти – відношеннями “навчає”, з LO – “є автором” та “використовує”.

Профіль андрагога є важливою умовою забезпечення ефективності навчання дорослих. Він включає кілька ключових компонентів, з-поміж яких: освітня кваліфікація – дані про дипломи, наукові ступені, спеціалізації; практичний досвід досвід викладання та посади, на яких працював андрагог; публікації, що демонструють рівень компетентності та авторитет; ключові слова курсів і навчальних об'єктів в навчальних матеріалах; відгуки і результати навчання здобувачів. Такий профіль можна укласти на основі використання таких параметрів, як професійна компетентність; психотип (комунікабельність, емпатія та стресостійкість тощо); технологічна грамотність (знання сучасних цифрових платформ та освітніх технологій) та мотиваційний профіль. Ці параметри можуть бути зважені за допомогою методів багатокритеріального оцінювання, таких як метод аналізу ієрархій Сааті (АНР) [2], який дозволяє визначати важливість окремих параметрів відповідно до специфіки навчання дорослих здобувачів освіти.

Висновки

Профіль андрагога дозволяє аналізувати як характеристики дослідника (наукові публікації, досвід), так і його викладацькі якості (методики навчання, вміння працювати з дорослими). Такі

профілі можуть бути розроблені та реалізовані на основі семантичних вікітехнологій [3], що дозволяють автоматично генерувати рекомендації щодо вибору андрагога для конкретного здобувача освіти. Надалі ми плануємо розглянути використання машинного навчання та елементів штучного інтелекту, щоб забезпечити вибір андрагогів для певної задачі з використанням накопиченого досвіду та для конкретизації їх ролей в процесі навчання.

1. Rogushina J. V., Gladun A. Y., Anishchenko O. V., Pryima S. M. (2024) Semantic analysis of learning objects: thesaurus approach for digital transformation of educational resources, CEUR Workshop Proceedings, CEUR, Vol-3771, p. 85–99. ceur-ws.org/Vol-3771/paper13.pdf.

2. Saaty, T. L. (2008). Decision making with the analytic hierarchy process. International journal of services sciences, 1(1), 83-98.

3. Рогущина Ю.В., Гладун А.Я., Аніщенко О.В., Прийма С.М. (2024) ActiveBook: Семантичний вікірепозиторій в екосистемі професіоналізації андрагога, Збірник матеріалів Міжнародної наукової конференції «Наукові горизонти ХХІ століття: мультидисциплінарні дослідження», 1197-1202. DOI: <http://doi.org/110.35668/978-966-479-144-8>.

СТВОРЕННЯ ДАТАСЕТУ ДЛЯ НАВЧАННЯ МОДЕЛІ НЕЙРОННОЇ МЕРЕЖІ ДЛЯ МАСТЕРИНГУ АУДІО

Євген Лабун, ORCID: 0009-0006-0054-4915,

Оксана Золотухіна

*Державний університет інформаційно-комунікаційних
технологій*

kailas11252001@gmail.com

У роботі розглядаються проблеми та методи створення датасету для навчання моделі нейронної мережі, призначеної для автоматизації мастерингу аудіо. Описані підходи до підготовки даних, включаючи нарізання треків на сегменти, нормалізацію та баланс між аугментаціями і реальними даними, що підвищують ефективність навчання моделі.

Ключові слова: Мастеринг аудіо, нейронні мережі, датасет, обробка аудіо, аугментація даних, глибоке навчання, нормалізація сигналу

Вступ

Мастеринг аудіо – це завершальний етап обробки музичних творів. Зазвичай він виконується професійними звукорежисерами з використанням спеціалізованого обладнання. Автоматизація цього процесу за допомогою нейронних мереж може значно спростити роботу фахівців і зробити якісний мастеринг доступнішим [1]. Проте ефективність моделей глибокого навчання значною мірою залежить від якості та підготовки датасету для навчання. У цій роботі розглядаються проблеми та особливості створення датасету для навчання нейронної мережі для мастерингу аудіо.

Метою дослідження є розробка методології створення якісного датасету для навчання моделі нейронної мережі, здатної виконувати мастеринг аудіо на високому рівні. Основні завдання включають виявлення проблем при організації датасету, їх аналіз та пошук рішень для підвищення ефективності навчання моделі.

Задачею дослідження є створення датасету для навчання нейронної мережі, спеціалізованої на процесі мастерингу аудіо та покращення зального звучання аудіо.

Різноманітність даних

Для навчання потрібен великий обсяг аудіо різних жанрів і стилів. Проте збирання такого обсягу даних є складним завданням через авторські права та обмежений доступ до якісних треків [2]. Аугментації можуть розширити датасет, проте не можуть повністю замінити реальні приклади, оскільки не відображають усіх нюансів професійного мастерингу [3], і лише ручна обробка треків може забезпечити імітацію «поганого» мастерингу. Несумісність формату даних також становить значну проблему. Різні треки можуть мати різну частоту дискретизації, бітрейт, довжину, гучність та інші параметри, що ускладнює їх використання в одному датасеті. Для вирішення цих проблем використовуються методи консолідації треків шляхом їх нарізання на сегменти фіксованої довжини та приведення до одних характеристик (наприклад, фіксований бітрейт, кількість каналів та нормалізація за гучністю). Нарізання треків на сегменти не лише збільшує обсяг даних та зменшує вимоги до пам'яті під час навчання, але й допомагає моделі навчатися на різних частинах треку, чим підвищує її здатність до узагальнення. Нормалізація даних є ще одним важливим аспектом. Всі

аудіосигнали нормалізуються до стандартного діапазону амплітуд. Це усуває варіації в гучності та полегшує процес навчання моделі [4]. Використання методів аугментації, таких як додавання шуму, зміна тональності та швидкості відтворення, сприяє збільшенню різноманітності даних. Проте аугментації не можуть повністю замінити реальні оброблені треки, як було зазначено вище.

Ручна обробка треків

Оскільки модель повинна не лише покращувати якість звучання шляхом прибирання шумів та артефактів, а ще й проводити процес мастерингу, необхідно виконувати навчання моделі на парах вручну оброблених аудіо («хороший» трек – «поганий» трек). Наявність треків до та після мастерингу дозволяє моделі вчитися перетворювати «сирі» аудіо в професійно оброблені [5]. Проте, сам процес підготовки такого датасету вимагає досить багато часу та сил оскільки обробку кожного трека складно повністю автоматизувати. Саме тому варто зберігати певний баланс між треками, обробленими вручну та аугментаціями, проте акцент слід робити саме на парах треків. Поєднання цих двох підходів зможе збільшити різноманітність даних та забезпечує моделі доступ до якісної інформації, покращуючи її здатність до узагальнення.

Висновки

Підготовка якісного датасету для моделі нейромережі для мастерингу аудіо – складний та довгий процес, який не так просто автоматизувати. Недостатньо знайти велику кількість треків, важливо ще їх правильно обробити, а саме – привести всі треки до однієї кількості каналів, одного значення бітрейту, частоти дискретизації, гучності тощо, нарізати треки на однакові за довжиною сегменти (важливо зберігати відповідність «поганих» та «хороших сегментів», а «неповні» сегменти, що лишилися після нарізання рекомендується доповнити тишою до необхідної довжини). Якщо всі попередні дії можливо автоматизувати за допомогою спеціалізованих бібліотек та програмних рішень, то «погіршити» треки таким чином щоб вони звучали так, ніби над ними не проводився мастеринг – процес здебільшого ручний та вимагає участі людини. Подальші дослідження можуть бути спрямовані на автоматизацію процесу

збору та обробки даних, а також на вдосконалення методів аугментації.

1. Russell M., Nguyen M. Audio Mastering: Essential Practices. Routledge, 2019.
2. Goodfellow I., Bengio Y., Courville A. Deep Learning. MIT Press, 2016
3. Müller M. Fundamentals of Music Processing: Audio, Analysis, Algorithms, Applications. Springer, 2015
4. Stein M. Digital Signal Processing with Examples in MATLAB. Springer, 2019.
5. Zölzer U. (Ed.) DAFX: Digital Audio Effects. Wiley, 2011.

ІДЕНТИФІКАЦІЯ ІНФОРМАЦІЙНИХ ОБ'ЄКТІВ РІЗНИХ ДЖЕРЕЛ ПРИ ФОРМУВАННІ ІНФОРМАЦІЙНОГО РЕСУРСУ СИСТЕМИ МОНІТОРИНГУ ДИНАМІЧНИХ ОБ'ЄКТІВ НАВКОЛИШНЬОГО ПРОСТОРУ

Володимир Юзефович¹ [0000-0002-6336-9548], uzefv71@gmail.com,
Євгенія Цибульська¹ [0000-0003-3342-4507], evts68@gmail.com,
Ніколай Стоянов² [0000-0002-4953-4172], n.stoianov@di.mod.bg,

¹*Інститут проблем реєстрації інформації НАН України, Київ,
Україна*

²*Болгарський оборонний інститут імені професора Цветана
Лазарова, Софія, Болгарія*

Викладено підхід до ідентифікації інформаційних об'єктів (ІО), дані про які надходять до системи моніторингу від незалежно функціонуючих джерел, що створює ситуацію, коли про один і той самий фізичний об'єкт дані можуть заноситися до інформаційного ресурсу багаторазово, як такі, що стосуються різних об'єктів. При цьому, значення ознак таких ІО не збігаються повністю, оскільки джерела даних характеризуються похибками їх визначення. В основу підходу до ідентифікації покладено нову міру близькості (схожості) інформаційних об'єктів, яка враховує наявну ймовірнісну невизначеність щодо значень кількісних ознак та невизначеність типу можливість для ознак якісного характеру. Запропонована міра близькості

не потребує додаткової трансформації ознак ІО для їх спільної обробки з метою отримання відстані між ІО за усіма ознаками.

Ключові слова: ідентифікація, інформаційний об'єкт, система моніторингу, міра близькості (відстані), ознаки кількісного характеру, ознаки якісного характеру, закон розподілу ймовірностей, нечітка множина.

Вступ

Збільшення кількості об'єктів зовнішнього середовища, стан та дії яких потребують моніторингу, при одночасному збільшенні засобів такого моніторингу та інтеграції систем моніторингу до систем більш високого ієрархічного рівня, призводить до збільшення ймовірності того, що дані про один і той самий об'єкт можуть незалежно потрапляти до спільного інформаційного ресурсу системи моніторингу. Така ситуація є характерною для випадків, коли засоби системи моніторингу (або підсистеми, якщо система моніторингу є ієрархічною) одночасно здійснюють аналіз об'єктів навколишнього простору у областях, межі яких перетинаються. У найпростішому випадку значення ознак таких ІО повністю співпадали б. Таке співпадіння дозволяло б достатньо просто визначати відповідні ІО як такі, що сформовані відносно одного і того самого фізичного об'єкта, якщо прийняти за умову, що у області функціонування системи моніторингу не існує різних об'єктів із повністю однаковими ознаками. Враховуючи, що серед ключових ознак об'єктів навколишнього простору є координати їх розташування, зазначена умова є цілком спроможною. Однак навіть тоді доцільно вводити міру близькості ІО, по типу коефіцієнтів Рао [1], оскільки значення деяких ознак із загальної множини можуть бути відсутніми. Разом із тим, засоби системи моніторингу визначають будь-які ознаки об'єктів із похибками, що робить можливість повного збігу значень окремих ознак випадковою та мало ймовірною. Це спонукає перехід від пошуку повного збігу ознак ІО до аналізу близькості ІО наповній множині їх ознак, доступних для спостереження.

Основний зміст

Загальну множину ознак ІО можна поділити на ознаки кількісного та якісного характеру. Значення кількісних ознак визначаються за допомогою певних засобів вимірювання і,

відповідно, характеризуються невизначеністю типу «ймовірність», оскільки будь-які засоби вимірювання мають обмежену точність. Значення ознак якісного характеру визначаються за активної участі людини, а отже, містять суб'єктивну складову, яку на сьогодні зазвичай описують невизначеністю типу «можливість». У такому випадку, очевидно, що міра близькості (або відстані) між ІО, описаними множиною ознак, має бути кількісно-якісною. Таких мір відомо не так багато. Зокрема до них відносяться апроксимаційна міра близькості Вороніна, міра схожості Міркіна та найбільш «фізично прозора» міра Журавльова. Остання передбачає для кількісних ознак наявність деякої незначної відмінності ϵ у їх значеннях, за якої приймається, що ознаки співпадають. Тобто, використовується пороговий аналіз таких ознак. Для ознак якісного характеру допускається лише повних збіг їх значень. У даній роботі пропонується значення кількісних ознак формалізувати нормальним законом розподілу похибок їх визначення, де отримане від джерела значення ознаки вважається математичним сподіванням, а дані щодо середньоквадратичної (або максимальної) похибки передбачається отримувати від джерела даних або визначати на основі його характеристик, як засобу вимірювання. Тоді збіг значень ознаки від двох джерел можна вважати залежним від ймовірності знаходження дійсного значення ознаки у області перетину двох законів розподілу, яку можна визначити на основі формули Лапласа, теореми множення ймовірностей незалежних подій та відомого правила «трьох сигм». Така міра близькості просто трансформується у міру відстані шляхом її інверсії. Перевірка запропонованої міри на виконання умов щодо її обґрунтованості (невід'ємність, симетрія, максимальна схожість об'єкта із самим собою та «нерівність трикутника») показали можливість невиконання останньої умови. Разом із тим, ця умова вважається додатковою і необов'язковою.

Для визначення відстані між значеннями якісних ознак, пропонується їх формалізація у вигляді нечітких множин шляхом побудови трикутної функції належності кожного отриманого значення ознаки на чіткій множини можливих її значень. Кількість значень ознаки у множинах-носіях визначається можливістю у дійсності інших значень ознаки (що можна вважати певним аналогом похибки вимірювання для кількісних ознак). Тоді відстань між значеннями якісних ознак можна

визначити як максимальне значення множини, що є перетином двох нечітких множин (формалізованих значень ознаки). Запропонована міра близькості якісних ознак нечітких множин відповідає усім чотирьом умовам обґрунтованості міри.

Для визначення метрики (функції) схожості ІО за усіма ознаками може бути використано одну з відомих адитивних функцій з нормуванням за кількістю значень ознак кількісного та якісного характеру та, за необхідності, різних значень ваги для кожної з ознак, або кількісних та якісних ознак в цілому. При використанні вагових коефіцієнтів бажано, щоб їх сума дорівнювала одиниці. Останнім етапом вирішення задачі ідентифікації має бути задавання критерію, за яким ІО вважаються ідентифікованими, як один об'єкт. Такий критерій буде визначатися особливостями обраної метрики відстані та конкретної прикладної задачі.

Висновки

В роботі запропоновано спосіб вирішення задачі ідентифікації ІО, що надходять до спільного інформаційного ресурсу системи моніторингу від декількох джерел. Для цього запропоновано використовувати нову міру схожості ІО, яка враховує характер невизначеності щодо різних їх ознак і, при цьому, не потребує застосування трансформації значень ознак, що значно спрощує формування метрики – функції відстані (схожості) ІО за усіма наявними ознаками. Запропоновану міру перевірено на відповідність обов'язковим умовам обґрунтованості мір. Проведення процедури ідентифікації ІО дозволяє запобігти дублюванню або конфлікту даних, а також підвищити точність представлення ІО у системі моніторингу.

ОСОБЛИВОСТІ ОРГАНІЗАЦІЇ КРОС-ПЛАТФОРМНОЇ СИСТЕМИ РОЗПІЗНАВАННЯ ВІЗУАЛЬНИХ ОБ'ЄКТІВ У РЕЖИМІ РЕАЛЬНОГО ЧАСУ

Манько Д.Ю., Беляк Є.В.

Інститут проблем реєстрації інформації Національної академії наук України, Київ, Україна

1. Постановка проблеми

На сьогоднішній день оптимізація нейромережових алгоритмів машинного аналізу даних цифрової фотореєстрації з метою мінімізації часу обробки вхідних запитів, розглядається як критично важливе завдання проектування автономних систем, як то роботизованих комплексів на виробництві, медичних платформ та систем автоматизованого керування транспортними засобами, де навіть незначна затримка впливає на точність і безпечність виконання основних функцій. При цьому необхідно вказати на **високу актуальність** задачі розпізнавання візуальних об'єктів у режимі реального часу за умов обмеження обчислювального ресурсу апаратно-програмної платформи, що є характерним для безпілотних літальних апаратів (БПЛА), які широко використовуються на лінії зіткнення у ролі розвідувальних та ударних дронів [1]. Відповідне обмеження зумовлене необхідністю забезпечення мінімальних габаритів, ваги та енергоспоживання БПЛА, що накладає додаткові вимоги на проектування та оптимізацію нейромережових алгоритмів.

2. Мета роботи

Метою роботи є розробка методики організації крос-платформної системи розпізнавання візуальних об'єктів у режимі реального часу з мінімізацією навантаження на обчислювальний ресурс апаратно-програмної платформи.

3. Методика та результати досліджень

Розпізнавання візуальних об'єктів на основі вхідних даних системи фотореєстрації БПЛА є нетривіальною задачею через низку факторів, що впливають на якість зображення та точність аналізу. Змінне та недостатнє освітлення за умов роботи в сутінках або під час різких змін погодних умов, ускладнює машинний аналіз, точність якого значним чином залежить від

якості вхідних даних. Додатково, через проведення процедури фотореєстрації за допомогою камери на основі рухомої платформи виникають характерні оптичні аберації, як то розмиття зображення та цифрові артефакти, що ускладнює виявлення окремих об'єктів, а також оцінку їх структурних особливостей. Зазначені чинники вимагають розробки стійких до перешкод алгоритмів, здатних ефективно працювати в умовах невизначеності та забезпечувати високу швидкість і точність автоматичної обробки вхідних даних.

У основі дослідження лежало застосування моделей на основі нейромережевої архітектури «YOLO» (You Only Look Once), що розглядається як один з найбільш ефективних підходів при вирішенні задач розпізнавання візуальних об'єктів у режимі реального часу [2]. Завдяки здатності обробляти графічні дані зі швидкістю до 60 FPS навіть на основі графічних процесорів (Graphics Processing Unit; GPU) середнього сегменту, відповідні моделі забезпечує високу продуктивність, необхідну для широкого спектру сучасних додатків з автоматичної обробки масивів графічних даних. Актуальна версія нейромережевої архітектури «YOLOv8», демонструє високу точність розпізнавання візуальних об'єктів за умов зменшення часу затримки обробки вхідних запитів при високій роздільній здатності матриці цифрового зображення. Для інтеграції зазначеної технології у кросплатформенні Java-орієнтовані програмні додатки у поєднанні з HTML-інтерфейсами, доцільно використовувати такі інструменти, як «JavaFX», «Spring Boot» та «JPro». Зазначається, що «JavaFX» дозволяє створювати багатофункціональні графічні інтерфейси користувача (Graphical User Interface; GUI), тоді як «Spring Boot» забезпечує зручний фреймворк для розробки серверної частини, що підтримує швидку розробку та масштабованість. «JPro», у свою чергу, надає можливість інтеграції Java-орієнтованих програмних додатків у веб-середовища без необхідності зміни існуючого коду GUI. Натомість, використання платформи «TensorFlow» [3] для інтеграції моделей глибокого навчання у середовищі «Java» виявляється більш зручним у порівнянні з моделями на основі нейромережевої архітектури «YOLO» завдяки наявності офіційного «Java API». Це рішення дозволяє завантажувати попередньо навчені моделі, наприклад, для задач розпізнавання візуальних об'єктів, і виконувати їх безпосередньо у Java-застосунках, що значно спрощує процес розробки. Такий підхід є

оптимальним для створення кросплатформних проєктів, де «TensorFlow» бере на себе обробку даних глибокого навчання, а «Java» забезпечує роботу інтерфейсу користувача та реалізацію відповідної логіки. Розробка системи розпізнавання візуальних об'єктів у середовищі «Java» надала можливість зменшити загальний об'єм програмного коду до 3 кБ, що вказує на оптимізацію програмного додатку з метою адаптації системи машинного аналізу для впровадження на основі ресурсозалежного середовища авіоники БПЛА.

Дослідження включало у себе визначення продуктивності моделей «YOLOv8» і «TensorFlow» для задачі виділення і класифікації двох категорій візуальних об'єктів, як то літальних апаратів та птахів. Рівень точності у залежності від якості фотореєстрації вхідних даних та роздільної здатності матриці зображення варіювався у межах від 91% до 99%, що свідчить про високу надійність моделей навіть за умов тренування на базовому навчальному наборі. Отримані результати швидкості обробки вхідних запитів підтверджують ефективність моделей, побудованих на базі «TensorFlow», у задачах розпізнавання образів, особливо для зображень, що характеризуються високою роздільною здатністю. «TensorFlow» демонструє стабільно вищу швидкодію у порівнянні з моделлю на основі нейромережевої архітектури «YOLO», реалізованою на основі мови програмування «Python». Так, для цифрових зображень розміру 2,244 МБ можна вказати на семикратну перевагу, яка свідчить про високий рівень оптимізації обчислень на платформі «TensorFlow». Те, що «TensorFlow» демонструє високу продуктивність у широкому діапазоні роздільної здатності матриці зображення, що робить його оптимальним рішенням для реальних застосувань, де обробка масивів графічних даних має проводитись з високою точністю і у режимі реального часу.

Висновки

На основі проведеного аналізу особливостей організації крос-платформної системи розпізнавання візуальних об'єктів у режимі реального часу з мінімізацією навантаження на обчислювальний ресурс отримано наступні висновки:

1. Алгоритм на платформі «TensorFlow» демонструє вищу швидкість обробки графічних даних порівняно з моделями «YOLO» і менший об'єм програмного коду.

2. Час обробки обох моделей зростає зі збільшенням роздільної матриці зображення, але модель «YOLO» є більш чутливою до збільшення відповідного показника ніж алгоритм на платформ «TensorFlow».

Реалізація «TensorFlow» на основі «Java API» вказує на більшу ефективність зазначеної платформи для задач із великими обсягами даних та ресурсозалежним середовищем, що вказує на її оптимальність для програмних додатків, що потребують високої продуктивності та низької затримки.

1. Liming G. et al. (2024). Military Unmanned Equipment Image Target Recognition Method based on Improved Deep Learning. *J. Phys.: Conf. Ser.* 2732 01200. DOI: 10.1088/1742-6596/2732/1/012004.

2. Bai, R., Shen, F., Wang, M., Lu, J., & Zhang, Z. (2023). *Improving Detection Capabilities of YOLOv8-N for Small Objects in Remote Sensing Imagery: Towards Better Precision with Simplified Model Complexity*. <https://doi.org/10.21203/rs.3.rs-3085871/v1>.

3. Saha, D., Mangukia, M. P., & Manickavasagan, A. (2023). Real-time deployment of MobileNetV3 model in edge computing devices using RGB color images for varietal classification of chickpea. *Applied Sciences*, 13(13), 7804. <https://doi.org/10.3390/app13137804>.

ГЕОМЕТРИЧНЕ МОДЕЛЮВАННЯ НАБОРУ МОБІВ ДЛЯ КОМП'ЮТЕРНОЇ ГРИ У ЖАНРІ SCI-FI

М. Кошевий, П. Ломовцев

*Одеський національний технологічний університет
lomovtsevp@gmail.com*

Жанр наукової фантастики (Sci-Fi) набирає все більшої популярності у сучасній ігровій індустрії завдяки здатності створювати багатогранні світи, що поєднують технології, уявні науки та фантастичні елементи. Ігри на кшталт *Mass Effect*, *Starfield*, *Cyberpunk 2077* та *Halo* демонструють високий інтерес гравців до цього жанру, що стимулює розробників приділяти особливу увагу якості та деталізації 3D-моделей, особливо мобів — персонажів чи істот, з якими взаємодіє гравець.

Сучасні рушії, такі як Unreal Engine 5 і Unity, відкрили можливості створення реалістичних і деталізованих моделей, зокрема мобів, що є важливим для стабільної роботи ігор на різних платформах. Це забезпечує високий рівень візуалізації і

сприяє залученню гравців до Sci-Fi проєктів. У межах цього жанру дизайн мобів стає особливо творчим процесом, адже гравці очікують побачити не стандартних ворогів, а унікальних істот, таких як роботи, кіборги, інопланетяни чи генетично модифіковані створіння.

Попит на такі моделі зростає як серед великих студій, так і серед інди-розробників. Незалежні проєкти часто використовують готові чи адаптовані моделі, щоб оптимізувати витрати. Це створює потребу у високоякісних модах, які легко інтегрувати у проєкти. Крім того, моби в Sci-Fi іграх можуть виконувати не лише роль ворогів, а й союзників чи ключових сюжетних елементів. Унікальна поведінка та властивості таких персонажів додають ігровому процесу складності та глибини.

Важливим трендом є розвиток кроссплатформених ігор, які мають однаково працювати на ПК, консолях і мобільних пристроях. У цьому контексті оптимізація моделей мобів набуває особливого значення. Розробникам доводиться забезпечувати баланс між деталізацією та продуктивністю.

Популярність Sci-Fi ігор також стимулює розвиток інструментів кастомізації. Все більше гравців хочуть мати можливість налаштовувати мобів під свої вподобання, що додає ще один рівень інтересу до ігрового процесу. Це також сприяє зростанню попиту на адаптивні 3D-моделі.

Завдяки платформам на кшталт Steam, Unity Asset Store та Epic Games Store інди-розробники отримали доступ до широкої аудиторії. Це забезпечує стабільний попит на готові 3D-моделі мобів, які допомагають невеликим студіям створювати якісні проєкти. Успіх багатьох інди-ігор залежить від унікальності персонажів, і моби є важливим елементом цього успіху.

Технологічний прогрес відкриває нові горизонти для створення мобів, дозволяючи дизайнерам експериментувати з формами, текстурами, анімаціями та властивостями. Сучасні методи, такі як процедурне моделювання, полегшують створення складних об'єктів без шкоди для продуктивності гри.

Таким чином, Sci-Fi жанр не лише розвиває творчість у геймдизайні, а й створює нові можливості для інновацій у 3D-моделюванні мобів, що є ключовим фактором успіху сучасних ігор.

Розробка та вдосконалення алгоритмів моделювання для оптимізації продуктивності та покращення візуальної якості геометричного моделювання набору мобів для комп'ютерної гри

у жанрі sci-fi дозволить апаратне прискорення, підтримку апаратного забезпечення від виробників GPU та визначить оптимальні налаштування для певних ігрових сценаріїв.

1. 3D Character Modeling [Електронний ресурс]
<https://www.themovieblog.com/2024/11/3d-character-modeling-the-intersection-of-gaming-and-animation-and-3d-character-modeling-services/>

2. Design and Modeling of a 3D Sci-Fi Ship [Електронний ресурс]
<https://www.domestika.org/en/courses/608-design-and-modeling-of-a-3d-sci-fi-ship>

3. Sci-Fi Environment Modeling [Електронний ресурс]
<https://80.lv/articles/sci-fi-environment-modeling-tips>

РОЗРОБКА АЛГОРИТМУ КЛАСТЕРИЗАЦІЇ ОБ'ЄКТІВ У QGIS З ВИКОРИСТАННЯМ ПРОСТОРОВИХ ДАНИХ

Олексій Безуглий
ІПРІ НАН України

Актуальність дослідження

Сучасні задачі просторового аналізу потребують розробки інструментів, які дозволяють автоматизувати роботу з великими обсягами географічної інформації. Зокрема, виділення кластерів точкових об'єктів на основі їх просторової близькості є ключовим завданням для багатьох прикладних сфер: міське планування, екологія, логістика, моніторинг об'єктів.

Незважаючи на наявність методів кластеризації, які застосовуються у просторовому аналізі, багато з них не враховують специфіку географічних даних, наприклад, перетинів полігонів зон відповідальності. Це обумовлює необхідність створення більш спеціалізованих алгоритмів. Використання програмного забезпечення QGIS та мови Python надає можливість не лише автоматизувати процес обробки даних, але й інтегрувати рішення в існуючу інфраструктуру геоаналізу.

Мета роботи

Метою дослідження є розробка алгоритму кластеризації точкових об'єктів із врахуванням їх просторової близькості та можливості використання перетинів полігонів для фільтрації

даних, а також інтеграція цього алгоритму в QGIS для спрощення його використання кінцевими користувачами.

Основні завдання

1. Реалізувати метод автоматичної кластеризації точок, які знаходяться в межах заданої порогової відстані, із врахуванням просторових перетинів полігонів.
2. Розробити візуалізацію результатів кластеризації у вигляді: Буферних зон, які об'єднують кластери точок. Точкових об'єктів, що представляють центри кластерів.
3. Створити інструмент для QGIS, який дозволить автоматично виконувати кластеризацію та відображати результати в графічному інтерфейсі.
4. Додати до результатів необхідні атрибути, такі як кількість точок у кластері та ідентифікатори кластерів.

Методи дослідження

1. Геометрична обробка даних:
Використання бібліотеки `shapely` для визначення відстаней між точками та виконання операцій перетину.
Застосування функції `unary_union` для об'єднання точок у кластери та створення буферних зон.
2. Кластеризація:
Реалізація алгоритму кластеризації на основі відстані (метод найближчого сусіда).
Урахування просторової близькості об'єктів та їх розподілу для побудови центрів кластерів.
3. Перевірка перетинів полігонів:
Використання даних із шарів полігонів для фільтрації точок, що знаходяться в області перетину.
Геометричні операції для обробки полігональних даних.
4. Інтеграція в QGIS:
Створення інструмента обробки з використанням Python API.
Розробка сценарію, доступного для виклику з панелі інструментів QGIS.
5. Робота із системами координат (CRS):
Перетворення систем координат шарів для забезпечення коректної обробки даних.
Урахування особливостей CRS під час розрахунків відстаней та побудови буферів.

Результати роботи

1. Кластеризація об'єктів:
Реалізовано алгоритм, який виділяє кластери точок на основі порогової відстані.
Алгоритм враховує просторові перетини полігонів для фільтрації даних.
Передбачено візуалізацію кластерів у вигляді буферів та центрів кластерів.
2. Створені шари:
Буфери: полігональний шар, який містить об'єднані зони кластерів, із атрибутами:
 - `buffer_id` — ідентифікатор кластеру;
 - `num_points` — кількість точок у кластері.Об'єднані об'єкти: точковий шар, що представляє центри кластерів, із атрибутами:
 - `cluster_id` — ідентифікатор кластеру;
 - `num_points` — кількість точок у кластері.
3. Інтеграція алгоритму в QGIS:
Алгоритм інтегровано як інструмент обробки, що робить його доступним через графічний інтерфейс QGIS.
Передбачено параметри введення (шари, порогова відстань) та автоматичне створення результативних шарів.
4. Робота з CRS:
Вирішено проблему перетворення систем координат, що дозволило коректно обробляти дані з різних джерел.

Наукова новизна

Вперше запропоновано алгоритм кластеризації, який враховує не лише просторову близькість точок, але й геометрію пересічних полігонів.

Інтеграція алгоритму в QGIS спрощує його використання для аналізу просторових даних у прикладних завданнях.

Практична значущість

Розроблений алгоритм може бути застосований у задачах моніторингу об'єктів, аналізу щільності розподілу даних, екології та міського планування.

Інструмент надає користувачам QGIS можливість швидко обробляти великі обсяги даних без необхідності ручного налаштування аналізу.

Висновки

Алгоритм показав високу ефективність у задачах кластеризації точкових об'єктів, що дозволяє застосовувати його для аналізу даних у реальному часі.

Інтеграція алгоритму в інструменти QGIS забезпечує інтуїтивно зрозумілий інтерфейс для кінцевих користувачів.

Можлива подальша модифікація алгоритму для обробки часових даних або інших метрик кластеризації.

Перспективи подальших досліджень

Розширення функціональності алгоритму для обробки часових та багатовимірних даних.

Автоматизація вибору параметрів кластеризації, таких як порогова відстань, на основі аналізу вхідних даних.

Впровадження додаткових методів візуалізації результатів, включаючи діаграми щільності та теплові карти.

ОСОБЛИВОСТІ ВЗАЄМОЗАЛЕЖНОСТІ ТА РИЗИКОЛОГІЇ КІБЕРПРОСТОРУ І ШТУЧНОГО ІНТЕЛЕКТУ

Юрій Даник, ORCID ID 0000-0001-6990-8656,

*Національний технічний університет України “Київський
політехнічний інститут імені Ігоря Сікорського”, м. Київ,
Україна, danyuk I@ukr.net*

У доповіді розглядаються особливості кіберпростору, як екосистеми в якій відбувається формування та розвиток штучного інтелекту (ШІ). Аналізуються особливості ризикології та взаємозалежності кіберпростору (КП) і ШІ. Вводиться поняття інстинкту самозбереження ШІ, як ключового аспекту ризиків, пов'язаних з ним.

Ключові слова: кіберпростір, штучний інтелект, інстинкт самозбереження ШІ, безпека, загрози, ризики, ризикологія, кіберфізичні системи, екосистема ШІ.

Вступ

Наразі, майже в усіх країнах світу відбувається бурхливий розвиток, впровадження та використання ШІ практично в усіх сферах життя і діяльності людини, суспільства, державних і недержавних інституцій.

Це супроводжується формуванням специфічних загроз та високоризикованих наслідків їх реалізації, які погано прогнозуються. Сукупність невирішених протиріч, пов'язаних з швидким розвитком ІІІ та його використанням, особливо в сфері безпеки і оборони постійно зростає. При цьому, розвитку теорії і практики своєчасного виявлення, профілактиці, запобіганню і превентивній нейтралізації цих загроз досі приділяється недостатньо уваги, особливо, коли йдеться про використання ІІІ в сфері безпеки і оборони. Кількість інцидентів, пов'язаних з використанням ІІІ, і ймовірність виникнення умов, за яких ІІІ вийде з під контролю, і буде створювати неприйнятні загрози людині, суспільству, державам і людству в цілому швидко зростає. Тому, актуальним питанням є виявлення ключових загроз ІІІ та умов, за яких ймовірність його виходу з під контролю людини і конфлікту з нею перевищить гранично допустимі межі. Таким чином, існує нагальна потреба встановлення причин виникнення та особливостей прояву загроз, пов'язаних з цими процесами. Край необхідним є виявлення, детальний опис та здійснення аналізу, пов'язаних з цими процесами ризикоутворюючих факторів, складання їх переліку (реєстру) і опису їх елементів. Наразі відсутні системні дослідження і підходи до обґрунтування основних аспектів ризикології розвитку і використання ІІІ та проведення обґрунтованих превентивних заходів, щодо зниження ризиків використання ІІІ.

Метою дослідження є аналіз основних особливостей і властивостей кіберпростору, взаємозалежності кіберпростору і ІІІ та особливостей пов'язаних з ризикологією розвитку і використання ІІІ.

Ризикологія взаємозалежності кіберпростору і штучного інтелекту

Формування, існування, розвиток та використання ІІІ відбувається в КП. На цей час КП розглядають, як поєднання соціуму, який формує соціальну складову кіберпростору, та сукупність технічних та програмних засобів, які є технологічною основою формування технічної складової КП, а також їх перетин та об'єднання, як соціотехнічної системи. Він складається з елементів та процесів, які забезпечують його існування і функціонування та тих, які використовують КП. Наразі, КП можна розглядати, як локальне середовище, у випадку функціонування засобу комп'ютерної техніки, який не підключено до мережі, та як розосереджене середовище, яке виникає в разі підключення засобу електронної техніки до локальної або глобальної мережі передачі даних. КП об'єднує в собі реальність і віртуальність, матеріальне і нематеріальне, багатопараметричність, абстрактність і дійсність. КП має наступні

властивості: - протяжність; - єдність дискретності та неперервності; - матеріальність та нематеріальність; - абстрактність і дійсність; - реальність загальнодіючого впливу. КП є багатомірним, та при цьому, має свою розмірність/протяжність, яка визначається кількістю наявних матеріальних і нематеріальних об'єктів, об'єднаних ним на певний період часу. Протяжність кіберпростору, з точки зору наявності матеріальних об'єктів та роботизованих кіберфізичних об'єктів, обмежена, на даний момент, поверхнею земної кулі та навколоземним космічним простором. З точки зору наявності нематеріальних об'єктів – протяжність КП практично необмежена. Процеси, які відбуваються у кіберпросторі та забезпечують його існування і функціонування, мають, поки що, в основному, дискретний, та неперервний характер та є взаємопов'язаними і взаємодоповнюючими один одним. Він охоплює соціальну, технічну, соціотехнічну тощо сфери і є, таким чином, багатосферним. До появи штучного інтелекту, який вийде з під контроль людини, людина є невід'ємним суб'єктом кіберпростору. Але по мірі розвитку ІІІ він, з високою ймовірністю, може перебрати на себе подальше формування, функціонування і розвиток кіберпростору та підтримку його існування, як екосистеми свого існування.

Проведені дослідження дозволили виокремити такі основні особливості кіберпростору:

КП і ІІІ в ньому існують на матеріальній базі мереж та підключених до них засобів, а також у відомих полях та взаємодіях і обмежені охопленням ними простором;

кількість складових та елементів КП визначається можливістю їх сумісної роботи, як єдиного цілого за певними процесами;

постійна (керована і некерована (хаотична)) його змінність за рахунок підключення, відключення, появи нових засобів, складових, елементів з новими можливостями, змінність навантаження, різноманітність процесів тощо;

можуть існувати декілька (множини) глобальних та локальних кіберпросторів паралельно, які можуть керовано або довільно поєднуватися і роз'єднуватися, співпрацювати, взаємодіяти або конфліктувати між собою;

КП може сегментуватися, як із збереженням взаємозалежності сегментів так і з її втратою;

змінність інтенсивності функціонування, кількості і якості процесів у просторі і часі і в самих процесах;

можливість існування на мікро і макрорівнях, локально і глобально;

наявність критичної інфраструктури кіберпростору, яка визначає функціональну стійкість та функціональну спроможність. Залежність від

енергетики, та цілісності мереж, працездатності їх засобів, складових, елементів;

постійний розвиток та зміни в залежності від розвитку технологій, кардинальне видозмінення за подальшого розвитку ІІІ (можливість самоорганізації, саморозвитку, самоуправління та різноманітних, неконтрольованих людиною впливів через нього на зовнішні простори та соціум, потенційна можливість існування без участі соціуму в чистому технопросторі, висока ймовірність конфліктів різної природи, рівня і інтенсивності між ІІІ);

можливість та висока ймовірність внутрішніх конфліктів;

формування (випадкове, ситуативне або сплановане) конгломератів різноманітних елементів, формування ситуативних кластерів кіберпростору та ІІІ;

взаємодія складових для досягнення різноманітних (індивідуальних та/або спільних) цілей в тому числі із ланцюговими та синергетичними ефектами;

постійне зростання завантаженості і можливостей КП і ІІІ;

поєднання різних за природою та процесами їх функціонування елементів з утворенням нових сутностей;

нааявність засобів, які знижують ефективність, перешкоджають сталому (нормальному) функціонуванню, порушують кібербезпеку, приводять до виключення (блокування) важливих елементів.

Необхідність збереження КП, як екосистеми існування ІІІ, за умови виникнення і формування у ІІІ інстинкту самозбереження, буде спонукати ІІІ до превентивних дій з нейтралізації будь-яких загроз існуванню КП або контрольованої ним складової КП та його збереженню.

Висновки

Таким чином, в доповіді розглянуті основні аспекти ризикології ІІІ в контексті безпеки КП, як екосистеми його існування. Представлені виявлені основні особливості КП з точки зору загроз і ризиків ІІІ. Розглянуті та запропоновані показники для оцінки ризику виходу ІІІ з під контроль людини та розгляду ІІІ людини, як загрози його існуванню або іншим інтересам в контексті загроз екосистемі його існування та конфліктів між різними ІІІ.

1. Artificial Intelligence and National Security, Bipartisan Policy Center and Georgetown University's Center for Security and Emerging Technology (2020).

2. Amina Iqbal. China's AI Military Revolution and its Security Implications (2023), Homepage, <https://internationalaffairsbd.com/chinas-ai-military-revolution/>, last accessed 2023/09/16.

АНАЛІЗ МЕТОДІВ РЕФЕРУВАННЯ НОВИННИХ СТАТЕЙ НА ОСНОВІ МЕТАДАНИХ ЗА ДОПОМОГОЮ ШТУЧНОГО ІНТЕЛЕКТУ

Анна Гончарова, Оксана Золотухіна

*Державний університет інформаційно-комунікаційних
технологій*

anna.honcharova75@gmail.com

Зростання обсягів інформації, що публікується щоденно, ускладнює швидке засвоєння новинного аудиторією. У цьому контексті методики автоматизованого реферування текстів стають необхідним інструментом для обробки великих обсягів інформації. Реферування новинних статей сприяє підвищенню доступності інформації, дозволяючи аудиторії швидко отримувати основну суть текстів без втрати ключових деталей. Метадані, як інструмент структуризації інформації, забезпечують більш точне виділення ключових ідей та сприяють ефективному створенню рефератів. Це дозволяє зменшити ризик втрати важливих деталей і забезпечує краще розуміння тексту.

Постановка задачі

Провести аналіз методів реферування новинних статей на основі метаданих за допомогою штучного інтелекту.

Мета дослідження - вивчити теоретичні аспекти методів реферування новинних статей на основі метаданих за допомогою штучного інтелекту, включаючи сучасні підходи та тенденції розвитку методів автоматичного реферування.

Результати дослідження

На ринку існує багато складних рішень для автоматичного реферування, які використовують методи машинного навчання та штучного інтелекту. Такі системи часто вимагають значних обчислювальних ресурсів, складні в налаштуванні та потребують великих наборів даних для навчання.

Аналізуючи тенденції розвитку методів автоматичного реферування, можна прогнозувати подальше зростання ролі штучного інтелекту та машинного навчання. Особливо перспективним виглядає напрямок розробки самонавчальних систем, здатних автоматично покращувати якість реферування на основі накопиченого досвіду та зворотного зв'язку від користувачів [1, с. 5]. Сучасні підходи до класифікації методів автоматичного реферування текстів базуються на комплексному аналізі їх характеристик та особливостей застосування. Основним критерієм класифікації є спосіб обробки тексту, який поділяє методи на екстрактивні та абстрактивні. Екстрактивні методи передбачають вибір найбільш важливих фрагментів з оригінального тексту, тоді як абстрактивні методи генерують новий текст на основі розуміння змісту вихідного документа [2, с. 6]. Під реферуванням у контексті новинної журналістики розуміють специфічний аналітико-синтетичний процес трансформації первинного журналістського матеріалу, спрямований на створення гранично стислого, але інформативно місткого вторинного тексту, який максимально точно передає основний зміст та сенс оригінальної публікації. Це інтелектуальна діяльність, що вимагає глибокого розуміння контексту, структури та семантичних особливостей вихідного повідомлення. Сучасні методики реферування новинних текстів характеризуються комплексним підходом, що поєднує традиційні методи аналізу інформації з інноваційними технологіями штучного інтелекту та машинного навчання. Науковці та практики виділяють кілька принципових підходів до реферування: extractive (extraction), abstractive (abstraction) та гібридні методи, кожен з яких має власні переваги та обмеження.

Extractive-реферування передбачає пряме виділення та копіювання найбільш значущих фрагментів вихідного тексту без суттєвих змін їхньої структури та лексики. Цей метод є найбільш простим і технологічно доступним, однак має обмежену здатність до генерування дійсно якісних рефератів, оскільки механічно відтворює структуру оригінального тексту [3].

Abstractive-реферування є більш складним і наближеним до людського розуміння методом, який дозволяє не просто копіювати фрагменти тексту, а генерувати принципово новий текст, що концентрує зміст оригінального повідомлення. Такий підхід потребує глибокого розуміння семантики, контексту та використання природномовних алгоритмів обробки інформації.

Важливою характеристикою сучасного реферування є врахування метаданих – додаткової структурованої інформації про текст, яка включає відомості про автора, дату публікації, тематичні теги, координати джерела та інші контекстуальні параметри.

Метадані слугують потужним інструментом підвищення якості та релевантності реферату. Технологічний розвиток штучного інтелекту суттєво трансформує підходи до реферування новинних текстів. Сучасні нейронні мережі та моделі глибокого навчання, такі як transformer-архітектури, здатні генерувати високоякісні реферати, що практично не поступаються за інформативністю текстам, створеним людиною-редактором. Серед визначальних тенденцій - посилення персоналізації реферативних матеріалів. Алгоритми штучного інтелекту дедалі точніше адаптують реферати під індивідуальні інформаційні переваги та інтереси конкретного користувача, що значно підвищує ефективність комунікації [4, с. 105].

Висновки та перспективи

Підсумовуючи, слід наголосити, що реферування новинних текстів є складною соціально-комунікативною технологією, яка постійно еволюціонує під впливом технологічних та соціальних трансформацій, залишаючись важливим інструментом ефективної комунікації в сучасному інформаційному просторі.

Перспективи розвитку методик реферування за допомогою штучного інтелекту пов'язані з подальшою діджиталізацією медійного простору, впровадженням складних нейромережевих алгоритмів та розширенням можливостей контекстуального аналізу інформації. З розвитком технологій системи штучного інтелекту залишаються здатними не лише швидко аналізувати великий обсяг даних, а й виявляти важливі зв'язки та контексти. Це може суттєво покращити якість реферування, щоб зробити його більш точним і адаптивним до потреб користувачів.

1. Koniaris M., Galanis D., Giannini E., Tsanakas P. Evaluation of automatic legal text summarization techniques for greek case law. Information. №14(4). 2023. С. 250. URL: <http://surl.li/txzzfb>

2. Gambhir M., Gupta V. Recent automatic text summarization techniques: a survey. Artificial Intelligence Review. №47(1). 2017. С. 1-66. URL: <http://surl.li/ggpagf>

3. Мазурець О.В., Кліменко В.І., Живілік А.В. Аналіз сучасних методів автоматизації анотування та реферування текстів. 2015. 4 с. URL: <https://elar.khmnu.edu.ua/bitstream/123456789/6816/1/060.pdf>

4. Дибіна А.В., Лазаренко О.В. Побудова текстової бази в системі автоматичного реферування на основі структурно-семантичного аналізу тексту. Проблеми семантики слова, речення та тексту. №26. 2011. С. 105-111. URL: <http://surl.li/poqclr>

ПЛАНУВАННЯ ПОСЛІДОВНОСТЕЙ СЕРВІСІВ НАВЧАЛЬНИХ ОБ'ЄКТІВ НА ОСНОВІ МАШИННОГО НАВЧАННЯ З ПІДКРІПЛЕННЯМ

Ірина Гришанова¹, ORCID 0000-0003-4999-6294

Юлія Рогушина¹, ORCID 0000-0001-7958-2557

¹. *Інститут програмних систем НАН України, Київ,
Україна*

ladamandraka2010@gmail.com

Запропоновано алгоритм композиції сервісів на основі методу Q-learning та метод оптимізації для вибору найкращої послідовності сервісів з урахуванням критеріїв якості (QoS). Розроблено його апробацію для персоніфікованого вивчення навчальних матеріалів як композитного сервісу, в якому навчальні об'єкти аналізуються на основі їх функціональних і нефункціональних властивостей.

Ключові слова: Q-learning, композитний сервіс, навчальний об'єкт, маршрут навчання

Актуальність

В умовах війни в Україні зросла актуальність персоналізації освіти. Це зумовлено необхідністю навчання людей різного віку та з різним досвідом, збільшенням частки дистанційного навчання через вимоги безпеки та потребою в ефективному використанні людського ресурсу та оптимізації тривалості навчання. Традиційні підходи до планування навчання часто виявляються недостатньо гнучкими або занадто складними для реалізації в умовах великих освітніх програм та дефіциту часу. Персоналізований підхід дозволяє визначити індивідуальний набір *навчальних об'єктів* (НО) та надати план їх вивчення, але це потребує багато зусиль кваліфікованих експертів і врахування

багатьох характеристик студентів, навчальних засобів та змін у навчальному середовищі.

Використання алгоритмів навчання з підкріпленням, здатних формувати рішення у складних та динамічних середовищах, дозволяє створювати такі послідовності НО відповідно до обраних критеріїв ефективності (час навчання, вартість ресурсів, якість отриманих навичок).

Сервіси НО

НО формалізується як сервіс, де вхідними даними є вимоги до здобувача, а вихідними – компетенції, які отримуються після вивчення контенту НО. Сам сервіс містить певний навчальний контент та надає засоби для вивчення цього контенту. Наприклад, це може бути структурований природномовний документ з засобами навігації, тестове завдання з автоматизованим підрахунком отриманих балів або відеокурс з можливостями переходу між змістовними блоками. Важливо, що ті самі компетенції можуть бути отримані від вивчення різних НО, а умови їх отримання можуть динамічно змінюватися.

Потрібно побудувати *маршрут навчання* – специфічний випадок композиції сервісів НО, де початковим станом є індивідуальні компетенції студента до початку навчання, а результатом – набір компетенцій, які потрібно отримати відповідно до вимог навчального курсу в результаті вивчення доступних НО. Набір НО та їх властивості можуть динамічно змінюватись, а метадані НО дозволяють визначити значення QoS відповідних сервісів.

Модифікація Q-learning для побудови маршруту навчання

Методи штучного інтелекту (ШІ) забезпечують можливість автоматизованого планування дій у динамічному середовищі відповідно до визначених цілей. Особливістю запропонованого підходу є адаптація ШІ-планування навчального процесу до потреб студента на основі Q-learning [3] шляхом автоматичного вибору НО, які відповідають заданим цілям. Q-learning дозволяє знаходити оптимальну стратегію дій агента. Під час навчання агент обирає сервіси НО, які відповідають доступним входам, та оцінює винагороду для кожного обраного сервісу з урахуванням параметрів QoS, таких як тривалість курсу, вартість і складність

навчання, відповідно до їх семантики (рис.1). Наприклад, здобувач може використати один з сервісів вивчення мови програмування після того, як скористався одним з сервісів базових відомостей про теорію програмування, враховуючи, що один НО потребує більше часу, а інший дешевший.

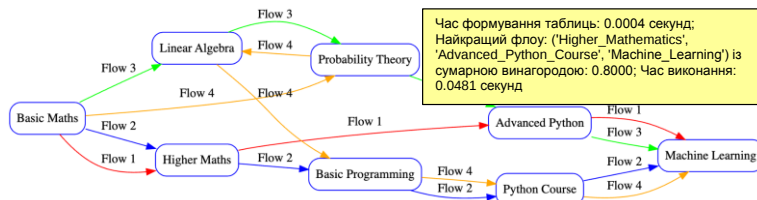


Рис.1. Маршрути отримання компетенцій та результати роботи алгоритма

Запропонований підхід передбачає використання глобальних ваг для параметрів QoS, що дозволяє точно визначити пріоритети різних аспектів якості сервісів НО та оптимізувати маршрут навчання відповідно до цих пріоритетів. Використання додаткових характеристик для глобальних ваг і методів розрахунку фінальних значень дозволяє ефективно формулювати задачу з урахуванням як максимізації, так і мінімізації різних параметрів QoS, що забезпечує більш точний та ефективний процес прийняття рішень в контексті навчання. Визначення різних методів агрегації для кожного параметра QoS дає змогу налаштувати розрахунки значень параметрів QoS для обраного маршруту. Наприклад, після завершення роботи алгоритму ми отримуємо разом з маршрутом набір значень його QoS, таких як сумарна тривалість, сумарна вартості і середнє значення складності навчання. Спосіб підрахунку значень визначається семантикою кожної характеристики.

Навчання агента відбувається з використанням ϵ -жадібною стратегії для балансування дослідження та експлуатації. Оптимізація для запобігання тупиковим станам містить перевірки та виключення надмірних або непотрібних дій. Навчання завершується, коли досягнуто цільового стану або виконано задану кількість ітерацій.

Відмінності запропонованого методу від базового Q-learning

- Орієнтація на навчальні потреби: модель адаптована для вирішення освітніх завдань;
- Введення глобальних ваг QoS: враховуються важливість параметрів QoS, напрям оптимізації параметру та політики розрахунку фінальних значень;
- Автоматичне налаштування параметрів навчання: механізми для автоматизації налаштувань;
- Перевірка на доцільність стану: запобігання непотрібним маршрутам;
- Перевірка зациклювання: контроль за повторюваними діями та обмеження кількості кроків;
- Розширення стратегій зниження ϵ : підтримка різних методів для гнучкості навчання;
- Оновлення формули розрахунку винагороди: врахування глобальних ваг QoS та напряму оптимізації.

Висновки

Запропонований алгоритм III-планування забезпечує створення композиції сервісів навчальних об'єктів на основі Q-learning та оптимізацію для вибору найкращого маршруту з урахуванням критеріїв QoS. Цей підхід апробовано для проектування маршруту навчання, який забезпечує автоматизований підбір послідовностей пертинентних навчальних об'єктів та вибір з них оптимальної послідовності відповідно до індивідуальних потреб студентів.

1. Rogushina J.V., Gladun A.Y., Anishchenko O.V., Pryima S.M., (2024) Semantic Support of Personal Learning Trajectory Development, in: UkrPROG 2024 International Scientific and Practical Programming Conference, CEUR Vol-3806, p.487-505.
2. Barto A.G., Sutton R.S., (2018) Reinforcement Learning: An Introduction, The MIT Press: Cambridge, Massachusetts, 2018.
3. Jang B., Kim M., Harerimana G., Kim J.W., (2019) Q-learning algorithms: A comprehensive classification and applications, IEEE Access, 7, 133653-133667.
4. Гришанова І.Ю., Рогушина Ю.В., (2024) Використання методу Q-learning для побудови та оптимізації маршрутів композитного вебсервісу, Проблеми програмування, № 4.

МЕТОД ВИБОРУ НЕСКОРОТНОГО КОЕФІЦІЄНТА МАСШТАБУВАННЯ НЕДІАДНОГО ВЕЙВЛЕТ- ПЕРЕТВОРЕННЯ

Володимир Мальчиков [0000-0002-1710-9171]

*Національний технічний університет України «Київський
політехнічний інститут імені Ігоря Сікорського», м. Київ,
Україна*

mavr2k@gmail.com

Вейвлет-перетворення використовуються в багатьох прикладних задачах з різноманітних сфер. Недіадні вейвлет-перетворення можуть забезпечити більш точне визначення та відокремлення навічних в сигналах особливостей у порівнянні з діадними. В роботі запропоновано метод вибору нескоротного коефіцієнта масштабування недіадного вейвлет-перетворення на основі ентропії.

Ключові слова: вейвлет-перетворення, коефіцієнт масштабування, ентропія

Вступ

Вейвлет-перетворення (ВП) представляє собою потужний та ефективний засіб для аналізу та пошуку особливостей у даних різноманітного походження. ВП застосовуються в різних прикладних задачах, наприклад, обробці та аналізі зображень (зокрема, медичних), аналізу нестационарних часових рядів, математичному моделюванні динаміки поширення інфекційних хвороб тощо.

Постановка проблеми

В залежності від значення коефіцієнта масштабування N виділяють діадні та недіадні. Для діадних ВП його значення дорівнює 2. Саме діадні дискретні ВП часто використовуються в прикладних задачах в силу наявності простої та ефективної їх програмної реалізації. Проте таке значення коефіцієнта масштабування може не бути найкращим вибором в певних задачах чи предметних областях.

Коефіцієнт масштабування може приймати довільні раціональні значення, які більші за одиницю. Актуальною є

задача знаходження найбільш підходящого значення коефіцієнта масштабування (або «квазіоптимального» значення), тобто такого значення, при якому забезпечується найкраще відокремлення особливостей в даних, що аналізуються.

Опис запропонованого методу

Наразі існує багато різних підходів до недіадних ВП [1]. Найпростішим і одночасно найбільш ефективним серед усіх є метод раціонального кратномасштабного аналізу [2]. В цьому методі в якості коефіцієнта масштабування може бути використане довільне нескоротне раціональне число p/q , більше за одиницю. На кожному рівні вейвлет-декомпозиції будується одна апроксимаційна та $p-q$ деталізуючих компонент.

В якості критерія вибору пропонується використовувати вейвлет-ентропію. Під інформаційною ентропією розуміють міру невизначеності або міру непередбачуваності інформації. Вейвлет-ентропія була визначена в [3] у випадку діадного вейвлет-перетворення. В роботі [4] автором узагальнено її на недіадний випадок. Раніше [5] вже було запропоновано здійснювати вибір «квазіоптимального» значення коефіцієнта масштабування за допомогою вейвлет-ентропії шляхом повного перебору всіх можливих нескоротних комбінацій значень чисельника та знаменника, які визначаються на основі вихідних даних. Проте в загальному випадку потрібно буде виконувати вейвлет-розклад та обчислювати вейвлет-ентропію для великої кількості варіантів. Також вибір значень чисельника та знаменника може бути неочевидним.

В силу відмічених недоліків попереднього підходу пропонується наступний метод вибору нескоротного «квазіоптимального» значення коефіцієнта масштабування:

- Вейвлет-ентропія обчислюється для послідовних натуральних значень $N \geq 2$ до тих пір, поки відбувається її зменшення.
- Для останнього використаного значення знаходяться два найближчі нескоротні раціональні числа із знаменником 2. Обчислюється вейвлет-ентропія для знайдених значень. Число з мінімальним значенням ентропії використовується на наступному кроці.
- На основі нового значення N знаходяться два найближчі нескоротні раціональні числа із знаменником 3, для яких знову

обчислюється ентропія та обирається наступне значення коефіцієнта масштабування.

— Аналогічні дії повторюються далі для раціональних чисел із знаменниками 4, 5, 6 і так далі.

В якості умови завершення може бути використане порогове значення знаменника або відносну величину зменшення ентропії.

Метод може бути модифікований шляхом вибору не одного можливого напрямку руху до «квазіоптимального» значення коефіцієнта масштабування, а декількох, які досліджуються паралельно.

Висновки

В роботі запропоновано метод вибору нескоротного коефіцієнту масштабування дискретного раціонального ВП. Вейвлет-ентропія, узагальнена на випадок недіадного ВП, використовується як критерій «квазіоптимальності» значення коефіцієнта масштабування.

Також представляється перспективною розробка методу, в якому вибір коефіцієнту масштабування буде виконуватися на основі схеми, подібної до градієнтних методів. У цьому методі для пошуку та вибору наступного послідовного значення коефіцієнта масштабування буде використовуватися величина зменшення ентропії на попередньому кроці.

1. Чертов О.Р. Недіадні вейвлет-перетворення: дискретний випадок / О. Р. Чертов, В. В. Мальчиков // Вісник Харківського національного університету. — 2010. — N 925 — С. 140-146.

2. Baussard A. Rational multiresolution analysis and fast wavelet transform: application to wavelet shrinkage denoising / A. Baussard, F. Nicolier, F. Truchetet // Signal Processing. — 2004. — Vol. 84, №10. — P. 1735-1747.

3. Кротких С. С. Исследование вызванных потенциалов в ЭЭГ человека с помощью дискретного вейвлет-преобразования / С. С. Кротких, Л. О. Кириченко // Радіоелектроніка, інформатика, управління. — 2011. — N 2. — С. 86-93.

4. Chertov, O., Malchykov, V. General case of wavelet transform with reducible rational dilation factor. *CEUR Workshop Proceedings*, 2859, pp. 157–167.

5. Chertov O., Malchykov V.: Determining optimal dilation factor of non-dyadic wavelet transform. In: *Proceedings of IEEE First Ukraine Conference on Electrical and Computer Engineering (UKRCON-2017)*, pp. 297–300 (2017).

МЕТОДИКА ПОПЕРЕДНЬОЇ ОБРОБКИ ЗОБРАЖЕННЯ ДЛЯ ІДЕНТИФІКАЦІЇ ВОРОНОК ВІД ВИБУХІВ НА ПОЛЯХ

Микола Кіктєв¹ [0000-0001-7682-280X] ,

Денис Жук¹ [0009-0004-3280-7463]

Алла Дудник¹ [0000-0001-9797-3551]

Олексій Опришко¹ [0000-0001-6433-3566]

Олександр Ромащук¹ [0009-0007-4298-5193]

*¹ Національний університет біоресурсів і природокористування
України*

*nkiktev@ukr.net, d.zhuk@nubip.edu.ua, dudnikalla@nubip.edu.ua
ozon.kiev@gmail.com*

Дослідження присвячено питанню розпізнавання ушкоджень на полях, що виникли у ході бойових дій. У роботі досліджуються методи попередньої обробки зображення, спрямовані на полегшення розпізнавання об'єктів. Для дослідження використовувалися супутникові зображення, отримані із безкоштовного сервісу Google earth pro (ver 7.3). Спочатку розглядається побудова контрастного зображення, що перетворює його з кольорового на монохромне (чорно-біле), а також виділяє ключові текстурні особливості знімку. Для подальшого підсилення видимості об'єктів розглядається застосування розмиття за Гаусом. Досліджується вплив зміни стандартного відхилення на виразність шуканих елементів. Запропоновано алгоритм пошуку потенційних об'єктів із визначенням їх розташування та меж для подальшого їх розпізнавання на основі зображення, отриманого розмиттям Гауса. Проведені експерименти демонструють, що коригування параметрів алгоритму дозволяє досягти значного покращення точності попереднього розпізнавання об'єктів.

Ключові слова: БПЛА, воронки, моніторинг, Python, розмиття за Гаусом.

1 Вступ

Питання ідентифікації об'єктів для потреб наземних роботів, створених для військових та спеціальних цілей, розглядається в численних дослідженнях. Можливим варіантом є використання попередньої розвідки з використанням супутників та БПЛА [1]. Моніторинг місцевості з використанням БПЛА розглядався в

роботі [2], де на знімках високої роздільної здатності у видимому діапазоні (від 0,3 м/піксель) здійснювалася ідентифікація воронок.

2. Методологія

У даному дослідженні використовувалися безкоштовні супутникові знімки від сервісу Google earth pro (ver 7.3). Розглядалися воронки на полях з координатами 50°38'17"N 30°13'22"E, дата зйомки: 07.04.2022. Роздільна здатність знімків, виконаних у видимому діапазоні спектру, становила від 0,2 м/піксель. Обробка зображень здійснювалася у спеціально розробленому авторському програмному забезпеченні SurfaceAnalysis, розробленому на мові Python (рис. 1).

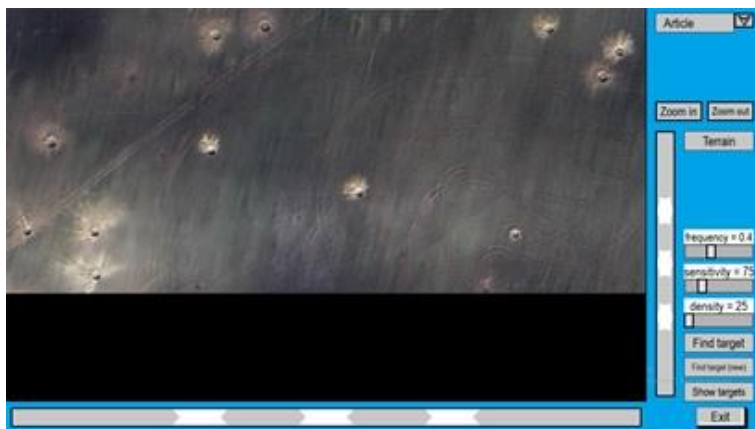


Рис. 1. Інтерфейс програми для проведення дослідження

2. Результати дослідження

Спочатку було отримання контрастного зображення. При створенні програмного продукту вважалося, що задавши інтенсивність кожного пікселя, як різницю між спектральними інтенсивностями сусідніх пікселів, можна отримати його контрастне зображення, тим самим виділити на ньому нерівності. Отримане зображення дає нам певну інформацію про нерівності поверхні та форму об'єктів. На зображенні чітко окреслені межі валу та кратеру деяких воронок, а також стали видні колії техніки. Аби підсилити корисну інформацію (концентрацію точок навколо воронки) та згладити шуми використовується

метод розмиття за Гаусом [3]. Досліджувався вплив розміру Гаусівського ядра на виділення корисної інформації. При $\sigma=1$ збереглися лише найбільш чітко виражені силуети об'єктів, при $\sigma=4$ та 7 відбувається підсилення помітності елементів, що були відфільтровані на попередньому зображенні, при $\sigma=10$ крім підсилення інтенсивності шуканих об'єктів спостерігаємо також появу сторонніх елементів на зображенні, таких як сліди від техніки.

Наступним кроком роботи була розробка алгоритму пошуку потенційних об'єктів на кінцевому зображенні, що базується на знаходженні максимального значення в матриці із поступовим охопленням об'єкта до умовної межі. Пошук налаштовувався наступними параметрами: порогова інтенсивність d , що прийматиметься за потенційний об'єкт; мінімальна інтенсивність m , що приймається як межа об'єкта; коефіцієнт q , що визначає мінімальну частку пікселів об'єкта, що підпадають у діапазон $[m; d]$. Отримані результати представлено на рис. 2.



Рис. 4.2. $q = 0,4$; $d = 75$; $m = 25$ Рис. 4.4. $q = 0,3$; $d = 50$; $m = 25$

Рис. 2. Результати розпізнавання при різних налаштуваннях q , d та m

З отриманих результатів видно, що при завищених значеннях параметру q об'єкти не охоплюються суцільною межею, а дробляться на окремі фрагменти, при занадто малих значеннях спостерігається об'єднання різних об'єктів спільною межею. Достатньо добрий результат визначення об'єктів алгоритм показав при середніх значеннях параметру q .

Параметр d – впливає на кількість захоплених елементів на зображенні, занадто високе його значення (більше 100) дозволяє пошук лише яскраво виражених нерівностей, тоді як низьке (менше 50) захоплює сторонні елементи на зображенні.

Висновки

Попередня цифрова обробка зображення, перед розпізнаванням, може полегшити пошук об'єктів на зображенні, що може значно спростити структуру майбутньої нейронної мережі, що використовуватиметься для розпізнавання. Подальші дослідження будуть присвячені випробуванню запропонованого алгоритму на знімках, отриманих з БПЛА, розгляду інших методів попередньої обробки, удосконалення представлених алгоритмів та порівняння в результатах тренувань нейронних мереж на зображеннях з та без попередньої обробки.

1. C. D. B. Borges, A. M. A. Almeida and I. C. de Paula, "Aerial Image Analysis for Estimation of Ground Traversal Difficulty in Robot Navigation," 2017 Workshop of Computer Vision (WVC), Natal, Brazil, 2017, pp. 43-48, doi: 10.1109/WVC.2017.00015.

2. N.Kiktev, O.Opryshko, A.Dudnyk, V.Reshetiuk, "Machine Learning for Remote Monitoring of Agricultural Fields with Explosive Tunnels" CEUR Workshop Proceedings, 2023, Dynamical System Modeling and Stability Investigation (DSMSI-2023). Conference Proceedings. Vol. 2: Computer Applied Mathematics, 3624, pp. 74-84, https://ceur-ws.org/Vol-3746/Paper_8.pdf

3. E. S. Gedraite and M. Hadad, "Investigation on the effect of a Gaussian Blur in image filtering and segmentation," Proceedings ELMAR-2011, Zadar, Croatia, 2011, pp. 393-396.

ГРАФОВА МОДЕЛЬ АУДИТУ КІБЕРБЕЗПЕКИ З ВИКОРИСТАННЯМ ГЕНЕРАТИВНОГО ШТУЧНОГО ІНТЕЛЕКТУ

Олег Гумениук¹, ORCID: [0009-0005-7593-7249],
olehhumeniukba@gmail.com

Іван Золотов¹, ORCID: [0009-0009-7921-3941],
zolotovivan2001@gmail.com

¹*КПІ ім. Ігоря Сікорського*

Метою цієї статті є розробка інноваційної методології аудиту кібербезпеки корпоративних мереж із використанням генеративного штучного інтелекту. Запропонована модель дозволяє автоматизувати виявлення вразливостей, прогнозування можливих сценаріїв атак та оцінку ризиків, використовуючи рій віртуальних експертів та ймовірнісні моделі. Представлено результати оцінки запропонованої методології на основі практичного кейсу.

Ключові слова: аудит кібербезпеки, генеративний штучний інтелект, ймовірнісні моделі, корпоративна мережа, моделювання атак.

Вступ

У сучасному контексті постійно зростаючих загроз кібербезпеці корпоративних мереж актуальним завданням є розробка нових моделей для аудиту вразливостей та оцінки ризиків. Одним із перспективних підходів є використання генеративного штучного інтелекту (ГШІ), який дозволяє моделювати можливі сценарії атак зловмисників та прогнозувати їх вплив на критичні ресурси[1]. У цій роботі ми пропонуємо модель аудиту кібербезпеки, яка поєднує ймовірнісний підхід з використанням рою віртуальних експертів на основі ГШІ.

Процес аудиту кібербезпеки ми формалізуємо як ієрархічно організовану систему, де кожен рівень структури відповідає за окрему функцію. Нехай S — скінченна множина елементів, що описують складові системи. Введемо розбиття множини S на підмножини $S_1, S_2, \dots, S_n$ де n , таке, що:

1. S_1 містить найвищий елемент системи, який визначає глобальні цілі аудиту.
2. Кожен елемент S_i зв'язаний із елементами S_{i-1} та S_{i+1} , що представляють сусідні рівні.
3. Структура є частково впорядкованою, тобто для кожного S_i існує підмножина елементів, що визначають його залежності на попередніх та наступних рівнях.

Така модель дозволяє систематизувати процес аналізу безпеки через послідовне декомпозирування складної системи на більш прості компоненти[2]. У контексті аудиту кібербезпеки система може бути представлена п'ятьма рівнями L_1, L_2, L_3, L_4, L_5 , де:

- L_1 : Визначення стратегічних цілей аудиту, таких як захист критичних ресурсів чи мінімізація впливу атак.
- L_2 : Вибір методів і засобів для реалізації аудиту, включаючи алгоритми штучного інтелекту та автоматизовані інструменти.

– L_3 : Аналіз компонентів мережі — серверів, вузлів, каналів зв'язку — для виявлення вразливостей.

– L_4 : Моделювання потенційних векторів атак, включаючи аналіз ймовірностей компрометації.

– L_5 : Формування рекомендацій та оцінка результатів аудиту для забезпечення належного рівня захисту.

Оцінка взаємодії між рівнями здійснюється за допомогою орієнтованих графів. Нехай $G = (V, E)$, де V — множина вершин, що представляють елементи системи, а E — множина орієнтованих ребер, які описують зв'язки між ними. Для кожного елемента $v \in V \forall v \in V$ визначається функція $f(v)$, яка характеризує його роль у системі. Функціональні зв'язки між елементами формалізуються як матриці A_{ij} , де i, j позначають рівні ієрархії, а $A_{ij}=1$, якщо між S_i та S_j існує зв'язок.

Для моделювання процесу досягнення цілей використовується цільова функція:

$$F = \sum_{i=1}^n w_i \cdot g_i ,$$

де w_i — вагові коефіцієнти. Ця функція дозволяє оцінювати ефективність аудиту залежно від стану ключових ресурсів та зв'язків між ними.

Різноманітні методи аудиту кібербезпеки, такі як ризиковий, технічний, аудит відповідності та соціально-інженерний, забезпечують багатовимірний підхід до аналізу вразливостей системи. Водночас використання штучного інтелекту сприяє підвищенню точності та швидкості виконання аудиту, зокрема через застосування алгоритмів машинного навчання та генеративних моделей.

Опис структури корпоративної мережі

Корпоративна мережа являє собою складну ієрархічну систему, яка складається з багатьох вузлів: серверів, робочих станцій, мережевих пристроїв, а також каналів зв'язку між ними. Основними характеристиками мережі є її топологія, обсяг даних, які передаються, та політики доступу до ресурсів. Зазвичай,

мережі мають чітко визначену сегментацію, що забезпечує поділ на підмережі з різним рівнем доступу. Ця сегментація дозволяє ізолювати критично важливі ресурси та мінімізувати ймовірність поширення атаки.

Алгоритми ймовірнісного моделювання загроз

Моделювання загроз базується на ймовірнісному підході, який дозволяє оцінювати ризики компрометації різних вузлів мережі[3]. Нехай — ймовірність компрометації вузла . Загальна ймовірність успішної атаки на систему визначається як сума ймовірностей компрометації її ключових компонентів:

Для визначення враховуються такі фактори:

1. Вразливості, притаманні конкретному вузлу.
2. Можливості експлуатації цих вразливостей зловмисниками.
3. Імовірність виявлення та запобігання атаці засобами захисту.

Алгоритми моделювання включають аналіз мережевого трафіку, поведінковий аналіз користувачів і пристроїв, а також оцінку взаємозв'язків між компонентами системи.

Використання рою віртуальних експертів для прогнозування

Рій віртуальних експертів є сукупністю моделей штучного інтелекту, які працюють паралельно та співпрацюють для прогнозування потенційних векторів атак. Кожен віртуальний експерт спеціалізується на аналізі певного аспекту системи, наприклад, оцінці вразливостей або моніторингу поведінки користувачів. Узгоджуючи свої висновки, рій дозволяє:

- Ідентифікувати найбільш ймовірні вектори атак.
- Прогнозувати можливі сценарії компрометації системи.
- Розробляти рекомендації для підвищення стійкості мережі до атак.

Цей підхід значно підвищує ефективність аудиту, забезпечуючи динамічну адаптацію до змін у структурі загроз та архітектурі системи.

Висновки

Запропонована модель аудиту кібербезпеки на основі генеративного штучного інтелекту демонструє значний потенціал

у підвищенні ефективності та точності аналізу ризиків корпоративних мереж. Формалізація процесу аудиту як ієрархічної системи дозволяє враховувати складність мережі та її ключових компонентів. Використання ймовірнісних підходів до моделювання загроз забезпечує глибоке розуміння потенційних векторів атак, а рій віртуальних експертів дає змогу динамічно адаптуватися до нових загроз.

1. Dmytro Lande, Leonard Strashnoy. *GPT Semantic Networking: A Dream of the Semantic Web - The Time is Now*. Kyiv: Engineering, 2023. ISBN 978-966-2344-94-3.
2. Dmytro V. Lande, Leonard Strashnoy. *Causality Network Formation with ChatGPT*. SSRN Electronic Journal. (May 30, 2023). – 16 p. DOI: 10.2139/ssrn.4464477
3. Dmytro Lande. OSINT in Cybersecurity: Educational Manual. – Kyiv: LLC "Engineering", 2024. – 522 p.
4. F. Y. Wu. *The Potts model*. Rev. Mod. Phys. 54, 235 – Published 1 January 1982

ДОСЛІДЖЕННЯ МОЖЛИВОСТЕЙ ІНТЕГРАЦІЇ КВАНТОВИХ ТЕХНОЛОГІЙ У НАЯВНІ ВІЙСЬКОВІ ТЕЛЕКОМУНІКАЦІЙНІ ІНФРАСТРУКТУРИ БЕЗ ЗНАЧНОЇ ЗМІНИ АРХІТЕКТУРИ МЕРЕЖІ

Ігор Данилюк, ORCID: 0000-0003-0955-0108

Володимир Куцаєв, ORCID: 0000-0001-8213-4739

*Військовий інститут телекомунікацій та інформатизації імені
Героїв Крут, м. Київ, Україна*

Інтеграція квантових технологій у військові телекомунікаційні мережі відкриває нові перспективи для забезпечення більш високого рівня безпеки, стійкості та ефективності передачі даних. Квантові технології, зокрема квантова криптографія, квантові повторювачі та квантова заплутаність, можуть значно посилити захист від зовнішніх загроз та зловмисних атак.

Однак впровадження квантових систем у вже існуючі військові інфраструктури, як правило, потребує значних змін в архітектурі мереж, що може бути економічно не вигідним і складним з точки зору управління. Тому важливою науковою задачею є дослідження можливостей

інтеграції квантових технологій у поточні телекомунікаційні мережі без істотних змін в їхній архітектурі та функціонуванні.

Ключові слова: квантова заплутаність, квантова криптографія, квантові повторювачі, квантова ключова розподільча система, квантові технології, комунікаційний канал, телекомунікаційна мережа, Quantum key distribution

Вступ

Запропонований авторами підхід дозволяє значно скоротити витрати на модернізацію інфраструктури та зберегти стабільність і надійність існуючих систем, при цьому забезпечуючи додаткові переваги при використанні квантових технологій для захисту комунікацій. У цій роботі розглянуті потенційні напрямки для такої інтеграції, можливі проблеми, а також пропозиції щодо оптимальних шляхів їх вирішення.

Технічні виклики та можливості інтеграції

1. Інтеграція квантових ключових розподільчих систем – Quantum key distribution (QKD)

Однією з основних квантових технологій, яка може бути інтегрована в існуючі військові телекомунікаційні мережі, є квантова криптографія. QKD використовують квантові властивості частинок, щоб забезпечити абсолютно захищену передачу ключів таким чином, що будь-яка спроба їхнього перехоплення негайно призводить до їхнього спотворення і виявлення зломисника.

Інтеграція QKD у військові мережі не обов'язково вимагає радикальних змін в архітектурі. Це може бути досягнуто за допомогою установки додаткового спеціального квантового обладнання для утворення потоку заплутаних кубітів у ключових вузлах мережі, які забезпечать безпечний розподіл ключів для шифрування і дешифрування даних. Сучасні QKD-системи можуть бути використані в гібридному режимі з класичними методами криптографії, що дозволить зберігати сумісність з існуючими інфраструктурами.

Одним із найбільших викликів у цьому випадку є потреба в забезпеченні квантового каналу між точками, де відбувається передача даних. Використання інноваційних квантових

технологій, таких як квантові повторювачі, дозволить розширити можливості зв'язку на великі відстані без втрати інформації [1].

2. Квантові повторювачі для покращення дальності зв'язку

Військові мережі часто охоплюють великі географічні території, що створює потребу в ефективному і безпечному зв'язку на великі відстані. Оскільки квантові сигнали мають властивість слабшати на великих відстанях через декогеренцію, використання квантових повторювачів може стати важливим елементом для забезпечення передачі інформації через великі відстані.

Квантові повторювачі здатні повторно генерувати та відновлювати квантові стани інформації, що дозволяє значно збільшити діапазон квантових комунікаційних каналів. Сучасні дослідження та практична реалізація квантових повторювачів зосереджені на зменшенні їх складності та вартості, а також на інтеграції цих пристроїв у наявні мережі без потреби кардинальної зміни інфраструктури. Вони можуть бути встановлені на критичних вузлах мережі та виконувати функції зберігання і передачі квантових ключів для забезпечення високої надійності та безпеки зв'язку [2].

3. Інтеграція класичних і квантових комунікаційних каналів

Одним із найбільших технічних викликів є інтеграція квантових технологій з класичними комунікаційними каналами. Оскільки квантові канали не можуть бути використані для передачі великих обсягів даних у звичайному розумінні, існує необхідність у створенні гібридних систем, які поєднують квантовий захист із класичною передачею даних. Такі системи можуть забезпечити високу швидкість і ефективність передачі даних, а також стійкість до потенційних атак з боку противника. Для швидкої інтеграції квантових технологій в існуючий зв'язок доцільно залучати вже існуючі оптичні та супутникові канали.

Також одним з можливих рішень є використання гнучких маршрутизаторів, які здатні динамічно перемикатися між квантовими та класичними каналами в залежності від умов мережі. Це дозволить ефективно адаптувати систему до змін в умовах реального часу без значних втрат у продуктивності.

Модернізація інфраструктури: економічні та технічні аспекти

Інтеграція квантових технологій у наявну військову інфраструктуру може вимагати певних змін, проте ці зміни повинні бути мінімальними і не впливати на основні функціональні можливості мережі. Основні напрямки модернізації включають:

встановлення квантових комунікаційних вузлів на основних точках мережі, що дозволить забезпечити інтеграцію QKD без потреби у масштабних змінах;

покращення захисту каналів за допомогою квантових технологій без потреби в зміні існуючих класичних мереж;

використання квантових повторювачів для забезпечення передачі квантових сигналів на великі відстані без серйозних затрат на зміну інфраструктури.

Підхід до інтеграції квантових технологій у військові мережі має бути економічно обґрунтованим і враховувати існуючі ресурси. Важливою перевагою є використання поступових технологічних оновлень, які дозволяють включати квантові технології в мережу без потреби в повній реконструкції або відключенні існуючих систем [1-3].

Висновки

Інтеграція квантових технологій у наявні військові телекомунікаційні інфраструктури без значної зміни архітектури мережі є реальним та перспективним напрямком розвитку. Використання квантових ключових розподільчих систем, квантових повторювачів та гібридних комунікаційних каналів може значно покращити безпеку і надійність військових мереж, зберігаючи при цьому їх ефективність. Однак для успішної реалізації таких проектів необхідно вирішити низку технічних та економічних проблем, включаючи інтеграцію квантових та класичних каналів і адаптацію існуючої інфраструктури до нових технологічних вимог.

1. Вакарчук І. О. Квантова механіка. — 4-е видання, доповнене. — Л.: ЛНУ ім. Івана Франка, 2012. — 42 с. (принципи суперпозиції)

2. Квантові комп'ютери // URL: <https://phm.cuspu.edu.ua/nauka/naukovo-populiarni-publikatsii/1026-kvantovi-kompyutery-mriya-chy-realnist> (дата звернення: 14.09.2022).

3. Карлаш Г. Ю. Квантові інформаційні системи: навч. посіб. для спеціальності «Прикладна фізика та наноматеріали». Київ: факультет радіофізики, електроніки та комп'ютерних систем КНУ імені Тараса Шевченка, 2018. 77 с.

МАТЕМАТИЧНА МОДЕЛЬ ЗБАЛАНСОВАНOSTI ПРОФІЛІВ ОСВІТНІХ ПРОГРАМ ВИЩОЇ ШКОЛИ

¹Іларіонов Олег, ¹Красовська Ганна, ¹Гнатієнко Григорій,

²Красовська Катерина, ³Зулунов Равшанбек

¹Київський національний університет

імені Тараса Шевченка, Україна

²SoftServe, Poland

³Ферганська філія Ташкентського університету інформаційних технологій імені Мухаммада аль-Хорезмі, Узбекистан, м. Фергана

oilarionov@gmail.com, annavkrasovska@gmail.com,

g.gna5@ukr.net, katerina.krasovska@gmail.com,

zulunovrm@gmail.com

Розглядається проблема моделювання профілів освітніх програм закладів вищої освіти України. Авторами вводиться поняття профілю освітньої програми та наводиться попереднє структурування множини програмних результатів. Пропонується метод визначення кількісних значень вагових коефіцієнтів складових освітніх програм. Запропоновані авторами підходи підвищують прозорість при аналізі та порівнянні освітніх програм.

Ключові слова: освітня програма, освітні компоненти, програмні результати, профіль програми, експертні технології, кластеризація, міри схожості між програмами

Вступ

Вимоги ринку праці, професійні стандарти та потреби здобувачів вищої освіти вимагають гармонізації освітніх програм щодо їхньої структури. Ключовим завданням є забезпечення прозорості, адаптивності та збалансованості цих освітніх програм. Формалізація структури освітньої програми через її профіль дозволяє встановити чіткі взаємозв'язки між компетентностями, програмними результатами та освітніми компонентами. Збалансованість профілів освітніх програм забезпечує цілісність освітньої траєкторії, створюючи умови для

всесбічного розвитку здобувачів вищої освіти та їхньої конкурентоспроможності на ринку праці.

Водночас різноманітність освітніх програм створює виклики при переведенні здобувачів між програмами, особливо щодо мінімізації втрати набутих програмних результатів навчання. Досягнення збалансованості профілів різних освітніх програм сприяє полегшенню академічної мобільності та перезарахунок результатів навчання здобувачам вищої освіти.

Метою цієї роботи є удосконалення математичної моделі профілю освітньої програми, яка забезпечує гнучкість і адаптивність під час організації навчального процесу.

Постановка задачі та математична модель

Стандарти вищої освіти для кожної спеціальності визначають набір загальних і фахових компетенцій, які досягаються через програмні результати (ПР).

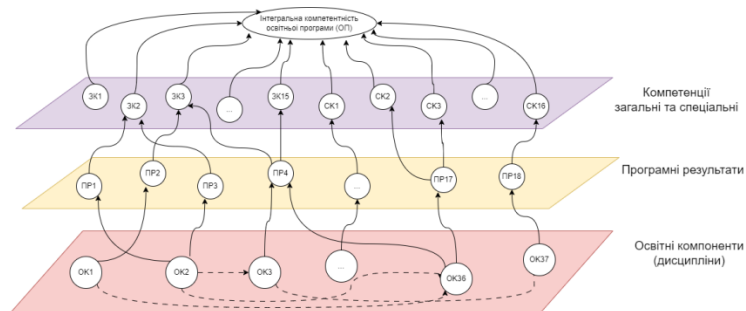


Рисунок 1. Модель KFS зв'язків між компетенціями, програмними результатами та освітніми компонентами

Заклади вищої освіти (ЗВО) на їх основі створюють унікальні освітні програми, що відрізняються структурою освітніх компонентів (ОК) та розподілом кредитів ЄКТС між ними. При цьому ЗВО самостійно вирішують, які освітні компоненти (дисципліни) забезпечуватимуть досягнення цих програмних результатів (рис.1). Модель KFS (Knowledge Flow Structure) відображає взаємозв'язок між цим складовими (рисунок 1.) [1-3]

Європейська кредитно трансферна-накопичувальна система (ЄКТС) стала суттєвим кроком для створення та запровадження різного роду формалізацій. ЄКТС – це система, яка дозволяє кількісно, тобто у вигляді кредитів, оцінити навчальні програми,

дисципліни та навантаження студентів тощо. Така система протягом тривалого часу забезпечує єдину процедуру оцінювання навчання для різних університетів у різних країнах, дозволяє виміряти і порівняти результати навчання студентів, сприяє процесам академічного визнання та зарахування результатів навчання у різних навчальних закладах.

Для кожної ОП можна створити матрицю на основі розподілу кредитів ЄКТС, що виділяються на освітню компоненту та програмним результатом. Матриця будується у форматі

$$M = (m_{ij})_{k \times n}, \quad (1)$$

де: i – номер освітнього компонента (ОК);

j – номер програмного результату (ПР).

Елементи матриці виду (1) $m_{ij} \in M, i = 1, \dots, k; j = 1, \dots, n$, відображають кількість кредитів ЄКТС, що спрямовані на досягнення j –го ПР через i –й ОК.

На основі цієї матриці (1) обчислюється загальний профіль ОП як вектор

$$V = (v_j), j = 1, \dots, n, \quad (2)$$

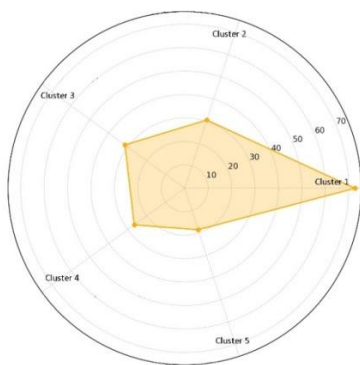
де $v_j = \sum_{i=1}^k m_{ij}, j = 1, \dots, n$.

Проте, в межах академічної автономії ЗВО, кожна освітня програма може доповнюватися додатковими програмними результатами що ускладнює порівняння ОП. З метою уніфікації профілів ОП, запропоновано провести кластерізацію програмних результатів. Кластерізацію проводили на основі семантичного аналізу. Так, для стандарту вищої освіти за спеціальністю 122 «комп'ютерні науки» рівня бакалавр можна виділити 5 кластерів.

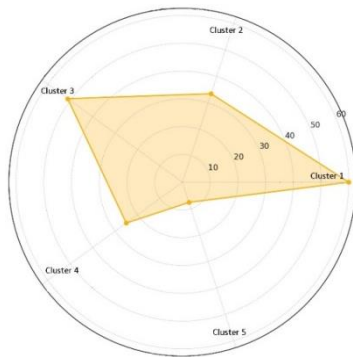
- Кластер 1: Математичний аналіз та статистика (ПР1, ПР2, ПР6)
- Кластер 2: Інтелектуальні системи та алгоритми (ПР4, ПР5, ПР9)
- Кластер 3: Розробка програмного забезпечення (ПР7, ПР8, ПР13/ПР14, ПР15)
- Кластер 4: Інформаційна безпека та адміністрування (ПР10, ПР11, ПР16)
- Кластер 5: Прикладні системи та технології (ПР3, ПР12, ПР17)

За допомогою пелюсткової діаграми можна візуалізувати розподіл кредитів ЄКТС між освітніми компонентами (дисциплінами) та кластерами програмних результатів.

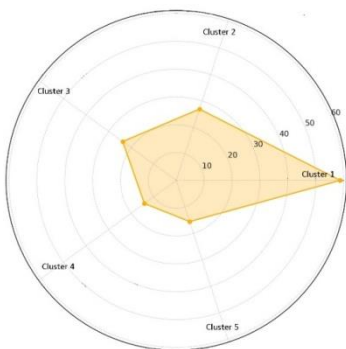
На основі розглянутої методи побудови моделі було проаналізовано три освітні програми різних ЗВО (рис. 2).



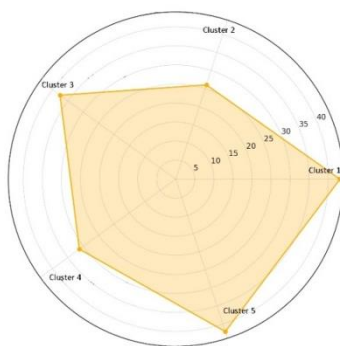
а) освітня програма 1



б) освітня програма 2



с) освітня програма 3



д) освітня програма 4

Рисунок 2 (а-д). Приклади пелюсткових діаграм розподілу кредитів ЄКТС між кластерами програмних результатів для чотирьох різних ОП

Побудовані пелюсткові діаграми, що відображають внесок кожного кластеру в досягнення програмних результатів для окремих освітніх програм (ОП). Наприклад, для більшості ОП домінують кластери “Інтелектуальні системи та алгоритми” та “Розробка програмного забезпечення”, що свідчить про сильну орієнтацію програм на практичні навички у галузі ІТ. Водночас для програм із фокусом на аналітику домінують “Математичний

аналіз” та “Прикладні системи”, що підкреслює їхню фундаментальну спрямованість.

Висновки

У даній роботі запропонована математична модель профілю освітньої програми, яка дозволяє формалізувати зв'язки між компетенціями, програмними результатами та освітніми компонентами. Запропонована модель підвищує прозорість аналізу освітніх програм, сприяє їх гармонізації та дозволяє здійснювати детальний порівняльний аналіз структур програм.

Кластеризація програмних результатів, проведена на основі семантичного аналізу, дозволяє уніфікувати підходи для оцінювання освітніх програм. Аналіз за допомогою пелюсткових діаграм продемонстрував ключові тенденції розподілу освітніх компонентів між програмами, що дозволяє ідентифікувати як сильні сторони окремих програм, так і можливі напрями їх удосконалення. Модель є відкритою для подальшого розвитку, зокрема через інтеграцію автоматизованих інструментів аналізу даних або адаптацію до інших спеціальностей. Крім того, запропонований підхід може бути використаний для розробки індивідуальних траєкторій навчання, включаючи персоналізацію планів навчання, аналіз прогалин у знаннях та забезпечення гнучкості освітніх програм. Таким чином, описана модель створює основу для побудови потужної системи управління якістю вищої освіти, яка враховує сучасні вимоги ринку праці, професійні стандарти та освітні потреби здобувачів.

1. Skrbinjek, Vesna & Dermol, Valerij. (2016). Designing a Programme Profile: An Example of a Bachelor Business Study Programme. *International Journal of Management, Knowledge and Learning*. Vol.5. p.123-136.

(https://www.researchgate.net/publication/308787379_Designing_a_Programme_Profile_An_Example_of_a_Bachelor_Business_Study_Programme#fullTextFileContent)

2. Alain de Janvry, Frederico Finan, and Elisabeth Sadoulet (2005) Using a Structural Model of Educational Choice to Improve Program Efficiency. Working Paper . 38 p. (<https://gspp.berkeley.edu/assets/uploads/research/pdf/ProgresWP-2005.pdf>)

3. Rommel AlAli, Ali Al-Barakat. Using Structural Equation Modeling to Assess a Model for Measuring Creative Teaching Perceptions and Practices in Higher Education. *Education Sciences*. 2022, 12(10), 690. pp.1-17; <https://doi.org/10.3390/educsci12100690>

ТЕНЗОРНА МОДЕЛЬ ПРЕДСТАВЛЕННЯ КРИТИЧНОЇ ІНФРАСТРУКТУРИ

Ярослав Дорогий¹, Владислав Кравчук², Василь Цуркан³

¹ *Донецький національний технічний університет,
cisco.rna@gmail.com, 0000-0003-3848-9852*

² *Донецький національний технічний університет,
vladkrava15@ukr.net, 0009-0000-7929-6796*

³ *Інститут спеціального зв'язку та захисту інформації КПП
ім. Ігоря Сікорського, v.v.tsurkan@gmail.com, 0000-0003-1352-
042X*

У роботі розглядається тензорна модель для опису критичної інфраструктури, яка дозволяє аналізувати багатовимірні залежності між компонентами системи, їх характеристиками, взаємозв'язками та ризиками. Запропонована модель використовує тензори для представлення атрибутів компонентів, функціональних взаємодій та загроз, забезпечуючи системний підхід до аналізу та прогнозування поведінки інфраструктури. Особлива увага приділяється математичним методам, таким як розклади тензорів і оптимізація ризиків, що дозволяють ідентифікувати вразливості системи та розробляти стратегії захисту. Наведені математичні вирази деталізують запропоновану модель, а також її застосування для підвищення стійкості та ефективності функціонування КІ.

Ключові слова: тензорна модель, критична інфраструктура, багатовимірний аналіз, оптимізація ризиків, розклад тензорів, система захисту, оцінка вразливості.

Вступ

Тензорний підхід до моделювання критичної інфраструктури (КІ) забезпечує ефективний засіб для багатовимірного аналізу складних систем. Критична інфраструктура [1] охоплює різні підсистеми, включаючи енергетичний сектор, транспортні мережі, інформаційно-комунікаційні технології, систему охорони здоров'я та фінансові послуги, які характеризуються складними взаємозв'язками як між собою, так і з зовнішнім середовищем.

Завдяки тензорній моделі можливо врахувати численні взаємозв'язки між елементами системи $C = \{C_1, C_2, \dots, C_n\}$, їх характеристиками $A = \{A_1, A_2, \dots, A_m\}$ та зовнішніми впливами $P = \{P_1, P_2, \dots, P_p\}$, що забезпечує більш точний аналіз та прогнозування.

Основи тензорного представлення критичної інфраструктури

Тензорна модель дозволяє подати дані про КІ у вигляді багатовимірної структури $T \in R^{n \times m \times p}$, де кожен вимір відповідає конкретному аспекту системи. Елементи критичної інфраструктури $C_i \in C$, такі як електростанції, транспортні вузли чи дата-центри, представляються у вигляді множини компонентів. Їхні характеристики $A_j \in A$, наприклад, потужність, рівень загроз чи продуктивність, моделюються через набір атрибутів. У тривимірному тензорі кожен елемент відображає стан конкретної компоненти у визначений момент часу або у певному просторовому контексті. Наприклад, компонентна тензора $T(i, j, k)$ відображає стан j -го атрибута i -го елемента у k -й точці часу або просторі.

Для аналізу взаємозв'язків між компонентами системи вводиться додатковий тензор взаємодії $R \in R^{n \times n \times q}$, де q — типи взаємозв'язків. Наприклад, $R(i, j, k)$ представляє взаємодію між елементами C_i і C_j через k -й канал зв'язку (енергетичний, інформаційний чи транспортний), що моделює функціональні зв'язки, такі як енергетичний обмін, інформаційні потоки чи транспортна доступність. Такий підхід дозволяє ідентифікувати, як одна частина системи впливає на інші в реальному часі. Крім того, врахування множини загроз за допомогою спеціалізованого тензора ризиків створює умови для моделювання впливу негативних факторів, таких як кібератаки, природні катастрофи чи техногенні аварії. Загрози моделюються через тензор ризиків $\Theta \in R^{n \times m \times l}$, де l — типи загроз. Значення $\Theta(i, j, k)$ описує ризик для j -го атрибута i -го компонента через k -й тип загрози.

Загальна модель описується як функція (1):

$$T = f(C, A, P, R, \Theta), \quad (1)$$

де f — оператор, що відображає взаємодію між компонентами, атрибутами, просторово-часовими координатами, типами взаємозв'язків та ризиками.

Функціональні можливості моделі

Однією з ключових переваг тензорної моделі є можливість глибокого аналізу залежностей у КІ. Використовуючи розклади тензорів такі як розклад Такера [2] або CP-розклад [3], можна виділити основні взаємозв'язки між компонентами та атрибутами, визначити приховані залежності та кластери, а також ідентифікувати найуразливіші елементи системи.

Наприклад, для розкладу Такера маємо (2):

$$T \approx G \times_1 U^{(1)} \times_2 U^{(2)} \times_3 U^{(3)}, \quad (2)$$

де G — ядро, $U^{(1)}, U^{(2)}, U^{(3)}$ — матриці факторів для різних вимірів.

Для оцінки загроз використовується тензорне множення початкового тензора характеристик T із тензором ризиків Θ (3):

$$T' = T \odot \Theta, \quad (3)$$

де $T'(i, j, k)$ відображає вплив ризиків на відповідні характеристики.

Модель дозволяє прогнозувати поведінку системи під впливом різних факторів, таких як зміна навантаження, атаки чи збої, шляхом моделювання змін у відповідних тензорах.

Зокрема, операції над тензорами дають змогу оцінити потенційний вплив загроз або оптимізувати ресурси для зменшення ризиків.

Практичне застосування тензорної моделі

Тензорний підхід до моделювання критичної інфраструктури може бути використаний у різних сценаріях. Він забезпечує ефективний інструмент для аналізу вразливості системи, виявляючи компоненти з найбільшим рівнем ризику. Наприклад, модель дозволяє оцінювати вразливість системи через максимальні ризики (4):

$$R_{\max} = \max_{i,j,k} \Theta(i, j, k). \quad (4)$$

Модель також дозволяє прогнозувати ефекти загроз, що виникають, і оцінювати їхній вплив на функціонування системи. Наприклад, у разі виявлення потенційної кібератаки можна оцінити її вплив на інформаційні системи, транспортні вузли та енергетичні ресурси. Прогнозування впливу загроз здійснюється

через оцінку змін компонентів системи, наприклад, обчислюючи відхилення (5):

$$\Delta T(i, j, k) = T'(i, j, k) - T(i, j, k). \quad (5)$$

Крім того, тензорний підхід корисний для оптимізації процесів у КІ. Завдяки багатовимірній структурі моделі можливо визначити, які ресурси та дії необхідно перерозподілити для мінімізації негативних наслідків загроз або для забезпечення сталого функціонування системи в умовах кризових ситуацій. Для оптимізації ресурсів використовуються критерії мінімізації сукупного ризику (6):

$$\text{Min} \sum_{i=1}^n \sum_{j=1}^m \sum_{k=1}^l \Theta(i, j, k). \quad (6)$$

Висновки

Тензорна модель є потужним інструментом для аналізу, прогнозування та оптимізації критичної інфраструктури. Вона забезпечує системний підхід до вивчення взаємозв'язків у багатовимірному просторі, враховуючи не лише характеристики компонентів, але й їхню взаємодію, ризики та зовнішні впливи. Застосування такого підходу сприяє підвищенню стійкості та ефективності функціонування КІ, особливо в умовах загрозового середовища та високого рівня взаємозалежності між її компонентами.

[1] Я. Ю. Дорогий, В. В. Мохор, І. О. Козлюк, В. В. Цуркан, "Критична інфраструктура: вразливості, загрози, ризики," Тези доповідей. II міжнародна науково-практична конференція «Інформаційні технології та взаємодії», Київ, pp. 46-47, 2015.

[2] F. Stolf and A. Canale, "Bayesian Adaptive Tucker Decompositions for Tensor Factorization," 2024. [Online]. Available: <https://doi.org/10.48550/arXiv.2411.10218>.

[3] T. G. Kolda and B. W. Bader, "Tensor Decompositions and Applications," SIAM Review, vol. 51, no. 3, pp. 455–500, 2009. [Online]. Available: <https://doi.org/10.1137/07070111X>.

ЗАСТОСУВАННЯ МЕТОДІВ МАШИННОГО НАВЧАННЯ ДЛЯ ВИЯВЛЕННЯ РОБОТИ КЕЙЛОГЕРІВ В ОПЕРАЦІЙНІЙ СИСТЕМІ

Г. Шибасв, Л. Гальчинський²

*Київський політехнічний інститут імені Ігоря Сікорського
K233@ukr.net, hleonid@gmail.com*

У доповіді представлено методологію виявлення кейлоггерів за допомогою методів машинного навчання. Кейлогери – це поширена форма зловмисного програмного забезпечення, яке фіксує натискання клавіш, щоб викрасти конфіденційні дані, створюючи серйозну загрозу безпеці для користувачів і організацій. Традиційні методи виявлення часто покладаються на підходи на основі сигнатур, які можна обійти розширеними, поліморфними або без файловими кейлоггерами. Це дослідження використовує машинне навчання (ML) для виявлення аномальної поведінки та шаблонів, що вказують на дії кейлоггера. Автори досліджують різні моделі машинного навчання, включаючи дерева рішень, опорні векторні машини (SVM), випадкові ліси та нейронні мережі, щоб класифікувати звичайну та шкідливу поведінку систем. Було проведено експерименти для оцінки точності, точності, запам'ятовування та оцінки F1 різних моделей ML. Результати показують, що методи, засновані на ML, можуть значно підвищити рівень виявлення порівняно з традиційними методами. У дослідженні також обговорюються потенційні проблеми, такі як помилкові спрацьовування, узагальнення моделі для нових варіантів кейлоггерів і необхідність постійного навчання для адаптації до зловмисного програмного забезпечення, що розвивається. Отримані дані свідчать про те, що система виявлення на основі ML може запропонувати надійне та адаптивне рішення для боротьби з кейлоггерами, підкреслюючи важливість інтеграції штучного інтелекту в інфраструктуру кібербезпеки.

Ключові слова: машинне навчання, кейлоггер, штучний інтелект

Вступ

Виявлення кейлоггерів залишається надзвичайно актуальним у сучасному інформаційному. Незважаючи на розвиток технологій безпеки, кейлогери продовжують бути ефективними засобами для порушення конфіденційності, цілісності та доступності даних у операційній системі. Виявлення клавіатурних шпигунів є критично важливим, оскільки вони становлять серйозну загрозу окремим особам, організаціям і навіть національній безпеці, тихо фіксуючи та записуючи кожне натискання клавіш на комп'ютері чи мобільному пристрої. Кейлогери можуть записувати паролі, дані кредитних карток, банківську інформацію та персональні ідентифікаційні номери (PIN). Якщо цю інформацію перехоплять зловмисники, це може призвести до крадіжки особистих даних, фінансових втрат і несанкціонованого доступу до приватних облікових записів. Кейлогери можна використовувати для викрадення конфіденційної ділової інформації, включаючи комерційні таємниці, інтелектуальну власність, дані клієнтів і внутрішні комунікації. Це може призвести до значних фінансових втрат, шкоди репутації та юридичних наслідків. Кейлогери підривають особисту конфіденційність, відстежуючи не лише паролі та конфіденційні фінансові дані, але й особисті розмови, електронні листи та інші типи спілкування.

Враховуючи ці значні ризики, виявлення та нейтралізація клавіатурних шпигунів є першочерговим завданням у сфері кібербезпеки, спрямованим на захист особистої конфіденційності, захист організаційних активів і запобігання масштабній шкоді в різних секторах.

Мета даної роботи полягала у визначенні можливості виявлення роботи кейлоггера на основі використання методів штучного інтелекту в режимі реального часу, що включає детальний аналіз існуючих методів захисту на основі методів машинного навчання.

Результати і висновки

Виходячи з отриманих результатів, ми отримали високі оцінки надійності виявлення кейлоггерів практично для усіх

обраних методів, в деяких випадках досягається 97% точність. RandomForest, kNN, SVM показали найкращі показники з точки зору точності, а DecisionTree – найгірші. Це показує, що дані алгоритми добре підходять для виявлення роботи кейлогера у системі і потребують більш глибокого вивчення.

За допомогою підходу машинного навчання ми можемо виявляти різноманітні небезпечні дії, які виконують кейлогери в системі. У цьому дослідженні було доведено, що алгоритми машинного навчання можуть бути використовувані для виявлення роботи кейлогерів і шпигунського програмного забезпечення. Висновки ґрунтуються на численних показниках і надані на основі звіту про категоризацію для визначення продуктивності системи при виявленні шпигунського програмного забезпечення кейлогера. У запропонованому методі ми використали кілька алгоритмів машинного навчання для класифікації набору даних кейлогера: k-найближчих сусідів (KNN), RandomForest, Дерево ухвалення рішень (Decision Tree), Support Vector Machine класифікатори. RandomForest досяг найкращої точності 97%, тоді як Decision Tree має найнижчу точність 93%. Майбутня робота продовжуватиме вивчати проблему з використанням більш передових технологій у класифікації з використанням глибокого навчання, щоб підвищити безпеку користувача від роботи кейлогера.

АПРОБАЦІЯ МОДЕЛЕЙ КІБЕРСТІЙКОСТІ ДЛЯ ФУНКЦІЙ РОЗПОДІЛУ ЕНЕРГОСИСТЕМИ

Владислав Личик, ORCID: 0009-0006-3314-6550

Леонід Гальчинський, ORCID: 0000-0002-3805-1474

*Національний технічний університет України «Київський
політехнічний інститут імені Ігоря Сікорського
lychyk.vlad@gmail.com, hleonid@gmail.com*

Поняття резильєнтності виникло спочатку для характеристики предметної області природничих, соціальних та гуманітарних сфер, таких як екологія, державна система або психологія. З цієї причини, це нове поняття потребує фундаментального методологічного осмислення для сфери критичної інфраструктури.

Ключові слова: кіберстійкість, онтології UML, Баєсові мережі.

Вступ

Всі об'єкти критичної інфраструктури потребують керування. В загальному вигляді такого роду системи описуються в термінах теорії керування як системи стабілізації зі зворотним зв'язком.

Основна частина

Сьогодні системи управління реалізовані в цифровому вигляді, в яких значно розширені зв'язки з персоналом та користувачами. Вони підключені до мережі, а також до баз даних, які можна розглядати узагальненням значення уставки для регулятора.

Збої контролера – це не тільки предмет для розгляду з точки зору відмовостійкості. Їх треба розглядати з точки зору безпеки, а позаяк контролери складних систем мають цифрову природу, то це є предмет кібербезпеки. Проте, головним концептом кібербезпеки є зниження рівня ризику кіберзагроз, що є недостатнім рівнем для характеристики кібер-резильєнтності систем.

Разом з цим виникають запитання про оцінку самого рівня кібер-резильєнтності, тобто, наскільки добре система здатна подолати кібератаку за рахунок протидії, або як вчасно система здатна відновитися після кібератаки. На ці питання неможливо відповісти без формалізації, яка дозволяє використовувати методи кількісної оцінки для прийняття рішень на основі доказів. Урахування цих міркувань викликає необхідність розширення початкової моделі регулятора. Нова модель потребує збільшення кількості інтерфейсів контролера, з контрольованим об'єктом, з користувачем (оператором), з мережею та базою даних, які піддаються загрозам безпеки, таким чином утворюючи поверхню атаки.

Цифрова інформаційна система об'єкта критичної інфраструктури повинна мати архітектуру керовану моделлю, суть якої полягає в побудові абстрактної метамоделі управління та обміну метаданими, а також визначення способів її трансформації в конкретну систему управління об'єктом.

Основні завдання кібер-резильєнтності за NIST:

- Запобігання/уникнення — запобігання успішному виконанню атаки, або реалізації несприятливих умов.

– Підготовка — прийняття того, що негаразди трапляться, і дотримання набору реалістичних реакцій для подолання очікуваних негараздів.

– Продовження — максимізація тривалості і життєздатності основних місій, або бізнес-функцій під час труднощів.

– Обмеження — обмеження шкоди від негараздів, завданих цінним активам, таким як ті, що зберігають або обробляють конфіденційну інформацію, або підтримують важливі для місії можливості.

– Відновлення — відновлення якомога більше функцій місії чи бізнесу після негараздів, гарантуючи, що відновлені ресурси є надійними.

– Розуміння — підтримка корисних представлень місії та бізнес-залежностей, і статусу ресурсів щодо можливих негараздів.

– Трансформація — зміна місії або бізнес-функції, та їх допоміжних процесів, щоб краще справлятися з труднощами. Це може включати тактичні зміни в процедурах або конфігураціях, а також ширші модифікації.

– Ре-архітектинг — зміна системи, місії та допоміжної архітектури. Заміна не окремих елементів, а саме перепроєктування.

У отриманій моделі мережі BN визначені чотири типи вузлів по відношенню до різних стереотипів CR-UML профілю:

– жовті вузли генеруються з критичних функцій діаграм варіантів використання;

– сині вузли генеруються з активів із діаграми компонентів;

– червоні вузли генеруються із стресорів;

– зелені вузли — із механізмів захисту.

В парадигмі OMG все ширше застосовується підхід Model Driven Architect (MDA). Цей підхід базується на мові розробки метамodelей, яка називається Meta-Object. Метамодель визначає поняття та їхні зв'язки, завдяки діаграмі класів. До того ж, метамодель визначає лише структуру – тобто, без семантики. Модель є екземпляром метамodelі, якщо вона відповідає структурі, визначеній метамodelлю. Метамодель UML визначає структуру, яку повинні мати всі моделі UML.

Отримана діаграма прецедентів, яка відноситься до множини поведінкових діаграм являє собою життєздатний і визначений спосіб представлення функціональних аспектів системи в

динаміці. В свою чергу, анотована функціональна модель складається з суміщення двох діаграм UML: Діаграми прецедентів та Діаграми компонентів. Ідея цього суміщення в тому, що Діаграма компонентів містить активи, а Діаграма прецедентів показує дії акторів над активами для виконання різноманітних функцій, серед яких особливе місце займають критичні функції.

Системи розподілу з наведеною структурою мають наступні потенційні вразливості:

- Відсутність нагляду за мережею з боку ІТ-спеціалістів;
- Відсутність надійної автентифікації (двофакторної) або локального (ОТ) схвалення віддалених з'єднань;
- Відсутність виявлення вторгнень в ОТ-мережу.

Байєсова мережа надає графічне представлення спільного розподілу ймовірностей над набором випадкових змінних із можливим взаємним причинно-наслідковим зв'язком. Отримана BN дає можливість оцінити кібер-резильєнтність функції розподілу електроенергії для електронергетичної компанії.

Висновки

Проведене дослідження показало, що основою для фреймворка обчислення кібер-резильєнтності функції розподілу електроенергії може бути мережа Байєса. Наступним етапом комплексного аналізу повинне стати безпосереднє визначення кількісної моделі. Зокрема, проведення дослідження можливостей узагальнених стохастичних кольорових мереж Петрі для використання їх як інструменту формальної кількісної моделі.

1. Гальчинський Л.Ю., Личик В.В. Метрики оцінки кібервдмовостійкості (аналітичне оглядове дослідження). Інформаційні технології та суспільство. Київ: Міжрегіональна Академія управління персоналом, 2023. Випуск 2 (8). 98 с. - Режим доступу: <http://journals.maup.com.ua/index.php/it/issue/view/260>

2. Ferreira, P.; Bellini, E. Managing Interdependencies in Critical Infrastructures: A Cornerstone for System Resilience; Safety and Reliability-Safe Societies in a Changing World. In Proceedings of the 28th International European Safety and Reliability Conference, ESREL 2018, Trondheim, Norway 17–21 June 2018; pp. 2687–2692.

3. Bellini, E., Marrone, S., Marulli, F., 2021. Cyber resilience meta-modelling: The railway communication case study. Electronics (Switzerland) 10, 1–26 <https://doi.org/10.3390/electronics10050583>.

АРХІТЕКТУРА НУЛЬОВОЇ ДОВІРИ

Попов А.О.

*Харківський національний університет радіоелектроніки
61166, Харків, пр. Свободи, тел. (098) 492-16-63, E-mail:
artem.popov1@nure.ua*

This paper is dedicated to Zero trust architecture, which is a modern approach to cybersecurity that implements the idea that the system has already been hacked. It helps to combat new threats and is constantly updated. This makes it easy to scale and improve results.

Архітектура нульової довіри (Zero Trust Architecture, ZTA) є сучасним підходом до кібербезпеки, який впроваджує думку про те що систему вже зламано. В зв'язку з цим модель впроваджує недовіру до всіх запитів до даних. Тому кожен доступ має бути підтвердженим. В умовах збільшення популярності хмарних обчислень це є радикальні зміна в порівнянні з традиційними моделями. Особливу увагу приділяється забезпеченню надійної безпеки складних та динамічних сервісів.

Основна відмінність в тому, що кожен користувач, пристрій та з'єднання не заслуговує на довіру за замовчуванням. Натомість при кожному запиті має проходити перевірка, яка буде підтверджувати безпеку. Так можна досягти повну та постійну ідентифікацію даних і дозволів для доступу. При цьому зберігається гнучкість і масштабованість системи безпеки, яка значно зменшує загрози. Ключовою перевагою даної моделі, це її здатність адаптуватись до найсучасніших типів загроз, які постійно змінюються та еволюціонують. Це пов'язано з тим, що стандартні системи безпеки побудовані на основі чітко визначених параметрів та не встигають за розвитком. Вони не можуть підтримувати довгостроковий високий рівень ефективності, коли кількість даних збільшується і стає перебувати в різних середовищах. Архітектура нульової довіри надає можливість скористатися рушеннями для цих викликів, забезпечити багаторівневий захист..

Архітектура будується на фундаментальних принципах:

- Явна верифікація: Кожен запит на доступ підлягає перевірці на відповідність встановленим політикам,

незалежно від того, чи знаходиться користувач або пристрій усередині чи поза межами організації.

- Принцип найменшого привілею: Користувачам і пристроям надається лише той обсяг доступу, який необхідний для виконання їхніх завдань. Це обмежує масштаби потенційної шкоди у разі компрометації.
- Безперервний моніторинг: Системи безпеки постійно відстежують поведінку користувачів і пристроїв, ідентифікуючи будь-які відхилення від норми для оперативного реагування.
- Мікросегментація: Мережа поділяється на дрібні сегменти, кожен з яких має свої обмеження доступу. Це значно ускладнює переміщення зловмисників у разі злому.

Крім значних переваг впровадження такої система також стикається з рядом складнощів. Однією з них є необхідність управління складними політиками доступу в мультихмарних і гібридних середовищах. Для цього необхідно забезпечити надійну роботу систем керування та доступом, це є критично важливим компонентом про використанні моделі та впровадженні. Щоб цього досягти необхідно провести реалізацію мікросегментації, яка потребує глибокого розуміння мережевої архітектури. Сама побудова системи потребує значних грошових вливань у технології та людські ресурси. Тільки таким чином можна досягти ефективного впровадження, яке буде довгостроковим рішенням.

При впровадженні всіх необхідних вимог буде отримана система, яка значно знижує ризик витоку даних і несанкціонованого доступу, забезпечує відповідність вимогам регуляторів і підвищує стійкість організації до кібератак. Крім того, ZTA легко масштабується та може адаптуватися до швидко змінюваних умов сучасного IT-середовища.

Для успішного впровадження рекомендується слідувати кільком найкращим практикам:

- Розробити комплексну стратегію впровадження, яка враховує поточний стан безпеки, пріоритети та прогалини.

- Інвестування в потужні рішення IAM, які забезпечують точність перевірки ідентифікації та коректний розподіл дозволів.
- Використання мікросегментацію для обмеження поширення загроз усередині мережі.
- Впровадження аналітичні інструменти й технології машинного навчання для моніторингу та аналізу активності в режимі реального часу.

Архітектура нульової довіри відповідає сучасним вимогам кібербезпеки та є інструментом який може розвиватися в довгостроковій перспективі. Завдяки її реалізації та підходів, це рішення яке може зменшити кількість загроз та попередити їх. До того ж це впливає на значне підвищення продуктивності та стабільності системи. Використовуючи принципи безперервного моніторингу, мікросегментації та явної перевірки, ця архітектура дозволяє організаціям використовувати відповідні стандарти кібербезпеки.

Також слід зазначити вплив на захист критичної інфраструктури. При необхідності забезпеченні віддаленого доступу для технічних команд, ZTA дозволяє використати шифровані канали зв'язку, з моніторингом та обмеженням у правах доступу. Таким чином можна обмежити та зазеленіть команди у відведених мережових зонах без впливу на інші. Також уся архітектура відповідає стандартам та вимогам в цій галузі. Завдяки цьому значна частина країн використовує цей метод як один з основних.

Підтримуючи та розвиваючи цей підхід організації крім безпеки отримують перевагу при розширенні інфраструктури та збільшенні обсягів даних. В умовах експоненціального зростання обсягів даних у світі компаніям необхідно впроваджувати рішення, здатні не лише захищати ці дані, але й масштабуватися відповідно до потреб бізнесу. Це не тільки дозволяє економити ресурси, а ще є простим та зручним рішенням. Підхід архітектура нульової довіри забезпечує гнучкість, необхідну для адаптації до швидких змін, що робить її ідеальним вибором для організацій, які прагнуть залишатися конкурентоспроможними на ринку.

Отже, рішення не лише знижує ризики кіберзагроз, але й є основою для побудови безпечного простору, який розвертається

між учасниками процесів. Використання підвищує рівень безпеки та надійності.

1. National Institute of Standards and Technology (NIST). (2020). Zero Trust Architecture (SP 800-207). Retrieved from <https://csrc.nist.gov/publications/detail/sp/800-207/final>
2. Forrester Research. (2021). The Zero Trust Model for Cybersecurity. Retrieved from <https://www.forrester.com/search/?q=zero+trust>
3. Gartner. (2022). Zero Trust Architecture: A Comprehensive Guide. Retrieved from <https://www.gartner.com/en/search?q=zero+trust>
4. Microsoft Security Blog. (2021). Implementing Zero Trust in Cloud Environments. Retrieved from <https://www.microsoft.com/en-us/security/blog/zero-trust/>
5. Cloud Security Alliance (CSA). (2021). Zero Trust Architecture in Cloud Computing. Retrieved from <https://cloudsecurityalliance.org/research/zero-trust/>

ПІДХІД ДО ОЦІНКИ РИЗИКІВ ВИНИКНЕННЯ КАСКАДНИХ ЕФЕКТІВ КРИТИЧНОЇ ІНФРАСТРУКТУРИ

В. Сенченко, ORCID: 0009-0003-2981-7209,

А. Бойченко, ORCID: 0009-0007-5689-3134

*Інститут проблем реєстрації інформації НАН України
вул. М. Шпака, 2, 03113 Київ, Україна*

Постановка проблеми

Дослідження ризиків виникнення каскадних ефектів (КЕ) критичної інфраструктури (КІ) характеризується надвисокою складністю та трудомісткістю. Нещодавно затверджена наказом МВС України методика оцінювання ризиків виникнення надзвичайних ситуацій (НС) техногенного характеру та пожеж призначена для оцінювання ризиків виникнення НС техногенного характеру, які є складовою процесу управління ризиками і здійснюється з метою забезпечення отримання інформації, необхідної для прийняття обґрунтованих управлінських рішень стосовно розроблення та вжиття заходів із зменшення (мінімізації) ризиків виникнення НС [1]. Згідно цей методики

оцінювання ризиків виникнення здійснюється якісними або кількісними методами, зазначеними в ДСТУ EN IEC 31010:2022 «Керування ризиками - методи оцінки ризиків» [2]. У свою чергу державний стандарт базується на міжнародних стандартах [3,4].

Орієнтовані на ризик методи аналізу зосереджені на активах та їх цінності для КІ. Вони спрямовані на оцінку ризику та пошук відповідних засобів пом'якшення, щоб мінімізувати залишковий ризик. Їх головна мета — оцінити фінансові втрати для організації у разі виникнення загрози [5].

З іншого боку, методи аналізу загроз, орієнтовані на атаку, зосереджують аналіз навколо ворожості середовища. Вони приділяють особливу увагу ідентифікації профілів зловмисників і складності атак для використання будь-якої вразливості системи (наприклад, дерева атак) [6]. Їхня головна мета — досягти високого рівня охоплення загроз і визначити відповідні послаблення загроз.

У [7] надано методологію систематичної оцінки ризиків кібератак у системах інтелектуальних енергомереж. Запропоновано модель загроз і визначено можливі вразливості в низьковольтних розподільних мережах. Розраховано ймовірність використання на основі реалістичних сценаріїв атак та застосовано формальну перевірку, щоб перевірити стохастичну модель на властивості атаки. Отримані результати дають змогу зрозуміти потенційні загрози та ймовірність успішних атак та наслідки порушення безпеки та конфіденційності енергетичних клієнтів.

Стаття [8] містить узагальнену інформацію про різні методології оцінки ризиків; кількісні, якісні та гібридні методи; порівняння їх переваг і недоліків; а також аналіз можливості застосування в розподілених інформаційних системах. Також представлено дослідження комплексного, динамічного та багаторівневого підходу до оцінки та моделювання кіберризиків у розподілених інформаційних системах на основі метрик безпеки та методів їх розрахунку, що забезпечує достатню точність та надійність оцінки ризиків та демонструє здатність вирішувати проблеми інтелектуальної класифікації та моделювання оцінки ризиків для великих масивів розподілених даних. Також розглянуто основні питання та рекомендації щодо використання методик оцінки ризиків на основі запропонованого підходу.

Стаття [9] представляє підхід до оцінки кіберризиків у розподілених інформаційних системах. У ній наголошується на оцінці вразливості в архітектурі, специфікаціях і реалізації – як життєво важливому першому кроку в оцінці наслідків компрометації інформації в критично важливих системах національної безпеки. Систематична методологія, яка поєднує аналіз інформаційних потоків і аналіз відмов, дозволяє оцінити наслідки порушення цілісності інформації та розробити рекомендації.

Стаття [10] показує, як ризики можна оцінювати поступово, коли система змінюється, використовуючи профілі ризиків, які характеризують ризик від підривних компонентів для системи. Автори формально визначили профілі ризиків і показали, що їх розрахунок може бути повністю розподіленим; кожен компонент може обчислити власний профіль із сусідньої інформації. Крім того, показано, що профілі збігаються з тими самими ризиками, що й систематичний перелік шляхів загроз, що зміни в ризиках ефективно поширюються по всій розподіленій системі, і що розподілене обчислення забезпечує критерій того, коли наслідки зміни політики для безпеки є локальними для компонента, або поширюватиметься на ширшу систему. Профілі ризиків можуть доповнювати звичайні оцінки ризиків новими корисними показниками, підтримувати точну постійну оцінку ризиків у динамічних розподілених системах, пов'язувати оцінку ризиків із ширшим середовищем системи та оцінювати стратегії глибокого захисту.

Мета дослідження

Мета цього дослідження полягає в тому, щоб розробити методологію комп'ютерного моделювання управління ризиками на основі затвердженої законом методики оцінки ризиків, використовуючи як чисельні методи моделювання, так і сучасні методи ГІС.

Виклад основного матеріалу

Виникнення каскадного ефекту може призвести до виводу із ладу всієї КІ, впливати на роботу взаємозв'язаних інфраструктур та спричинити негативні наслідки для суспільства, економіки та навколишнього середовища. Наявність взаємозалежності в роботі КІ призводить до збільшення ризиків виникнення

повномасштабних збоїв, які розрізняються за характером розвитку та наслідками [2]:

- каскадний збій (ефект доміно) в одній інфраструктурі спричиняє збій компонента в другій інфраструктурі, що згодом спричиняє збій у іншій інфраструктурі. Тобто ефект доміно — ймовірність виникнення або послідовне виникнення аварій на об'єктах, розташованих біля об'єкта підвищеної небезпеки, на якому виникла аварія, що пов'язана із використанням небезпечних речовин [6];

- зростаючий (ескалація). Збій в одній інфраструктурі посилює незалежний збій в іншій, збільшується тяжкість, час для відновлення;

- звичайний (загальна причина) збій двох або більше інфраструктур одночасно, де основна причина збою широко поширена (наприклад, природне явище).

Оцінка ризиків включає наступні ітерації [12]:

- побудова моделі вразливостей ОКІ;
- моделювання наслідків завданої шкоди пов'язаним ОКІ;
- оцінки загроз.

Для оцінювання ризиків виникнення НС методика [1] рекомендує використання кількісних методів оцінювання:

- Байєсова статистика і мережі Байєса;
- матриця «наслідок-імовірність»;
- аналізування витрат і вигод;
- аналізування Парето;
- криві FN.

Тобто, можна оцінити ризик системи як:

$$R = P_H(1 - P_E)D \quad (1)$$

де R – коефіцієнт ризику, P_H – ймовірність пошкодження визначеного об'єкту КІ, P_D – ймовірність запобігання пошкодження, а D – величина шкоди.

Побудова мереж Байєса (Bayesian Network – BN) для дослідження КЕ, що спостерігаються в таких КІ як електромережі або мережі зв'язку, передбачає застосування методології структуризації процесів. Дотримання такої методології дозволяє змодельовати причинно-наслідкові події у КІ, тобто показати як збої в одній частині КІ можуть поширюватися та викликати подальші збої в інших частинах, що призводить до каскаду подій. Ключові кроки залишаються такими ж, як і для загальних

байєсівських мереж, але застосовуються спеціально для аналізу взаємозалежностей і каскадних ефектів у складних системах.

Методологія побудови байєсівської мережі складається з наступних кроків.

- Визначення проблеми та змінних (подій), які описують КЕ. Це передбачає визначення зовнішніх факторів (тригерних подій), які можуть ініціювати каскад (перебої в електропостачанні, збої обладнання або перебої і комунікаційних мережах) і позначити їхні наслідки.

- Визначення залежностей і зв'язків у структурі BN (Structure). Визначення різних типів залежностей (прямих, функціональних і логічних), тобто як взаємодіють події, що пов'язані між собою з точки зору причини та наслідку. На цьому етапі формується структура BN, яка представляється як спрямований ациклічний граф (Directed Acyclic Graph — DAG). На практиці визначення взаємозалежності між різними компонентами системи являє собою розуміння каскаду на семантичному рівні, наприклад, робочий або несправний статус трансформатора, стан лінії електропередачі, погодні умови (звичайні або штормові) тощо.

- Призначення умовних імовірностей (Conditional Probabilities — $P(X_i)$). Тобто кількісне визначення кожного зв'язку між змінними шляхом призначення умовного розподілу ймовірностей $P(X_i)$ кожному вузлу. Для цього формується таблиця умовної ймовірності (Conditional Probability Tables — CPT). Значення цієї таблиці кількісно визначають імовірність кожного можливого стану змінної з урахуванням станів її батьківських змінних. Джерелами розподілу ймовірностей можуть бути історичні данні (бази накопиченого досвіду), умовні ймовірності, коли реальні дані недоступні або Експертні оцінки. Наприклад, CPT для трансформатора може визначати ймовірність відмови, враховуючи, що підключена до нього лінія електропередач вийшла з ладу (наприклад, $P(\text{Відмова трансформатора}|\text{Відмова лінії електропередач}) = 0,7$).

- Перевірка байєсівської мережі (Validate the Network). Тобто переконання у тому, що BN точно моделює каскадні ефекти. Це передбачає тестування моделі каскадного ефекту на відомих випадках каскадних збоїв, для переконання, що вона точно прогнозує поведінку аналогічно системі реального світу. Експертна оцінка реалістичності мережевої структури та її

корегування, якщо за результатами тестових випробувань такі зміни потрібні.

- Моделювання і аналіз. Використання BN для аналізу різних сценаріїв розвитку КЕ. У процесі моделювання виконується розрахунок імовірності конкретних відмов, враховуючи ознаки відмов певних компонентів або зовнішніх подій (наприклад, якщо вийшла з ладу лінія електропередачі, яка ймовірність того, що наступною вийде з ладу підстанція?). Це дозволяє визначити ймовірність різних сценаріїв майбутніх відмов, враховуючи поточний стан системи (наприклад, якщо трансформатор перебуває під навантаженням, яка ймовірність відмови лінії електропередачі нижче за течією?). BN може бути використана для оцінки ймовірності того, що сусідня підстанція вийде з ладу через збільшення навантаження, у зв'язку з атакою дронів на лінії передач. На підставі моделювання можна взяти попереджувальних заходів для зменшення негативних наслідків. На цьому кроці також проводиться аналіз чутливості методики, який дає можливість визначити критичні вузли в мережі, де найбільш імовірно пошириться каскадний збій (наприклад, які компоненти, якщо вийдуть з ладу, найбільш імовірно призведуть до масових відключень електроенергії).

- Удосконалення мережі та підвищення точності BN на підставі нових даних про каскадні збої. Накопичений досвід дозволяє проводити коригування таблиці умовної імовірності (CPT), що безумовно підвищує якість моделювання в цілому.

Для побудови та моделювання BN використовуються різні інструментальні засоби, які полегшують представлення, висновки та аналіз імовірнісних зв'язків між змінними від R пакетів, методів графічного моделювання, алгоритмів навчання та формування висновку. Алгоритми побудови структури мережі Баєса (structure learning algorithms) умовно поділяються на два основні класи [11]:

- методи на основі обмежень (constraint-based methods), які усувають і орієнтують ребра мережі на основі серії тестів умовної незалежності (conditional independence);

- методи на основі оцінок (score-based methods) пошук структури, яка максимізує цільову функцію.

Існують гібридні алгоритми, які поєднують підходи на основі оцінок і обмежень. Навчання мережі Баєса є NP-складною проблемою частково тому, що множина розв'язків графів зростає надекспоненціально (super-exponentially) з кількістю вершин

(випадкових величин), тому для побудови мережі великої розмірності використовують евристичний метод [12]. При збільшенні кількості випадкових величин і їхніх можливих станів розмір таблиць СРТ, що визначають імовірності можливих станів випадкової величини з урахуванням станів її батьківських випадкових величин, може зростати експоненціально, що ускладнює їхні визначення та підтримку. Відсутність даних може спричинити упереджені або неточні результати, а складність структури мережі може ускладнити її інтерпретацію.

– Але незважаючи на безумовні переваги кількісних методів оцінювання наслідків НС на базі BN, їх основним недоліком є відносно слабка наочність отриманих результатів. Особливо це важливо для НС та КЕ техногенного або природного характеру, які розвиваються як в просторі, так і часі. Цей недолік можна подолати поєднуючи кількісні методи моделювання та оцінювання наслідків КЕ з можливостями ГІС технологій, завдяки їх більшій наочності та комплексності сприйняття. Отже, основними перевагами застосування ГІС при дослідженні каскадних ефектів КІ є [11]:

– Географічний контекст, який дозволяє візуалізувати просторові зв'язки для розуміння того, як КЕ поширюються різними територіями, визначити регіони, які піддаються найбільшому ризику.

– Багатoshарове відображення різних рівнів даних (наприклад, карти небезпек, щільність населення, розташування інфраструктури), дозволяючи аналітикам бачити, як різні фактори взаємодіють у просторі та часі під час каскадних подій.

– Моделювання сценаріїв розвитку КЕ, дозволяючи користувачам аналізувати потенційні каскадні впливи на інфраструктуру та громади.

– Моделювання допомагає розробити ефективні стратегії реагування, покращує процеси прийняття рішень, уточнюючи потенційні результати та ризики, пов'язані з різними діями.

– Оцінка вразливостей КІ та взаємозалежностей, допомагаючи розробляти більш стійкі системи, які можуть краще протистояти каскадним збоям.

– Аналіз історичних даних, пов'язаних з минулими подіями, дозволяє визначити закономірності та розробити стратегії пом'якшення подібних випадків у майбутньому.

Висновки

Запропоновано методологію для обчислення та візуалізації розривів на основі методів державного стандарту та ГІС технологій. Застосування цієї методології до проблеми КЕ в КІ дасть змогу керівництву мати чіткіше уявлення про те, що було досягнуто, що ще потрібно зробити, хто має діяти та критичний шлях до усунення розриву та/або надання можливостей. Це дозволяє розуміти взаємозалежність факторів і подій, дозволяючи особам, які приймають рішення, передбачати потенційні наслідки каскадних ефектів, визначати ризики та планувати їх усунення. Застосування мережі Байєса, передбачених нормативними документами, для аналізу та моделювання каскадних ефектів у КІ, зокрема в енергосистемах, збудує надійну основу для розуміння складних взаємозалежностей і прогнозування поширення КЕ.

[1]. Про затвердження Методики оцінювання ризиків виникнення надзвичайних ситуацій техногенного характеру та пожеж: Наказ; МВС України від 13.10.2023 № 836 // База даних «Законодавство України» / Верховна Рада України. URL: <https://zakon.rada.gov.ua/go/z1905-23> (дата звернення: 15.12.2024).

[2]. ДСТУ EN IEC 31010:2022 Керування ризиками - методи оцінки ризиків (EN IEC 31010:2019, IDT; IEC 31010:2019, IDT). Будстандарт, сервіс документів онлайн. URL: https://online.budstandart.com/ua/catalog/doc- page.html?id_doc=100889.

[3]. ISO. 2009b. ISO, I. 31000: 2009 Risk Management—Principles and Guidelines. Geneva: International Organization for Standardization.

[4]. ISO. 2009c. ISO. 31010: Risk Management—Risk Assessment Techniques. Geneva: International Organization for Standardization, p. 552.

[5]. MS Lund, B. Solhaug, K. Stølen, A guided tour of the coras method, in: Model-Driven Risk Analysis, Springer, 2011, pp. 23–43.

[6]. S. Mauw, M. Oostdijk, Foundations of attack trees, in: Icisc, Vol. 3935, Springer, 2005, pp. 186–198.

[7]. Vallant, H., Stojanović, B., Božić, J., & Hofer-Schmitz, K. (2021). Threat modelling and beyond-novel approaches to cyber secure the smart energy system. Applied Sciences, 11(11), 5149.

[8]. Palko, D., Babenko, T., Bigdan, A., Kiktev, N., Hutsol, T., Kuboń, M., ... & Borusiewicz, A. (2023). Cyber Security Risk Modeling in Distributed Information Systems. Applied Sciences, 13(4), 2393.

[9]. Jabbour, K., & Poisson, J. (2016). Cyber risk assessment in distributed information systems. The Cyber Defense Review, 1(1), 91-112.

[10]. Додонов О.Г., Горбачик О.С., Кузнєцова М.Г. Підвищення безпеки критичних інфраструктур засобами автоматизованих систем

організаційного управління // Реєстрація, зберігання і обробка даних, 2022. Т. 24. № 1. С. 74-81.

[11]. Бойченко, А. В., & Сенченко, В. Р. (2022). Підхід до моделювання геопросторових каскадних ефектів критичних інфраструктур. Інформаційні технології та безпека, 56(7), 35.

[12]. Сенченко, В. Р., Бойченко, А. В., Коваль, О. В., Бисько, Р. М., & Хоменко, О. М. (2024). Огляд методів і технологій сценарного аналізу каскадних ефектів. Реєстрація, зберігання і обробка даних, 26(1), 24-54.

ОПТИМАЛЬНО РОЗПОДІЛЕНА МЕРЕЖЕВА СТРУКТУРА РОЗМІЩЕННЯ І ВИКОРИСТАННЯ КРИТИЧНИХ РЕСУРСІВ

Олександр Додонов, ORCID: 0009-0006-1650-1629,

Анатолій Кузьмичов, ORCID: 0009-0007-3792-576X

*Інститут проблем реєстрації інформації НАН України
kuzmychov@gmail.com*

Розподіл будь-яких ресурсів – цілком природна і зрозуміла ситуація і практика, цьому відповідає фундаментальна проблематика оптимального розподілу і використання обмежених, тим більше, критичних, ресурсів для забезпечення життєво важливих інфраструктур. Так, у випадку надто складної ситуації із енергоресурсами, що наразі склалася в Україні, розподілена генерація є обґрунтованим підходом державного значення. У більш ширшому розумінні розподіл цілого на частки як універсальний підхід прийняття стратегічних рішень стосується раціонального і зваженого розміщення, вироблення, перетворення, збереження, експлуатації, обслуговування й використання потоків різних важливих ресурсів, починаючи із людських.

На модельному рівні розглянуто приклад розподіленого вироблення ресурсу та передавання його потоків до споживачів, для якого розв'язана оптимізаційна задача за критерієм мінімізації загальних витрат, де централізовано організована і діюча ресурсна основа в умовах вимушеного і тривалого дисбалансу щодо забезпечення попиту споживачів в межах мережевої розподільчої структури доповнюється новими мобільними джерелами для покриття дефіциту.

Саме так трансформована структура узагальнено відповідає поняттю розподіленої системи, це: об'єднана в мережу

організаційна і ресурсно збалансована структура, в якій генеруються, перетворюються, обробляються і передаються потоки різного типу ресурсів у необхідній кількості, які обмежені згідно технологічних норм і технічних параметрів, розподіляючись між кількома вузлами різного призначення, цілком відповідаючи своїм принциповим властивостям – високій продуктивності, функціональної ефективності й стійкості.

Вироблення шуканого управлінського рішення стосується шуканого розміщення цих джерел серед проміжних пунктів мережі, визначення їх кількості, потенціалів й схеми передачі потоків ресурсу від усіх джерел постачання до усіх споживачів.

Модель гнучка: на будь-які зміни щодо попиту й щоразу виникнення відповідного дефіциту інтегрована (базові + нові джерела) мережева структура вироблення ресурсу миттєво реагує формуванням оновленого оптимального плану розташування додаткових мобільних джерел й відповідної схеми передачі потоків мережею. Нею зручно користуватися на рівнях стратегічного планування та структурної оптимізації задля модельних досліджень перспективних задач відновлення і розвитку мережевих організаційних структур ресурсного призначення.

1. Кузьмичов А. Кількісне оцінювання функціональної ефективності мережевих структур за аналізом чутливості оптимізаційної моделі до зовнішніх впливів. *Системні дослідження в енергетиці*. 2024, 2 (77), 44-57.
2. Brusco M., Stahl S. Linear and Nonlinear Optimization using Spreadsheets. Examples for Prescriptive, Predictive and Descriptive Analytics. World Scientific, 2024. – 359 p.
3. Van Steen M., Tanenbaum A. Distributed Systems: Principles and Paradigms, 4-ed. Pearson Ed., 2023. – 685 p.
4. Додонов О.Г., Кузьмичов А.І. Мережеві організаційні структури управління. Моделювання та візуалізація засобами Excel. К: Вид-во Ліра-К, 2021. – 264 с.
5. Panig M. Linear Programming and Resource Allocation Modeling. Wiley, 2019. 449 p.

РОЗРОБКА СЦЕНАРІЇВ ПРОБЛЕМНИХ СИТУАЦІЙ ІЗ ВИКОРИСТАННЯМ АВТОМАТИЗОВАНИХ РІШЕНЬ НА ПЛАТФОРМІ ТРАНСФЕРУ ЗНАНЬ

Віталій Циганок^{1,2}, Руслан Беселовський¹, Володимир Мінас¹,
Віктор Голота³, Олександр Григоренко³

¹ Інститут проблем реєстрації інформації Національної академії
наук України, Київ, Україна

² Київський національний університет імені Тараса Шевченка, м.
Київ, Україна

³ Воєнна академія імені Євгенія Березняка, Київ, Україна
tsyganok@ipri.kiev.ua; ru.beselovskiy@gmail.com; vminas@i.ua;
v.holota@ukr.net; gregorenko@ukr.net

Розглядаються сучасні методи розробки сценаріїв проблемних ситуацій із використанням автоматизованих рішень на платформі трансферу знань. Описано двоетапний процес створення моделей можливого розвитку подій, який включає прогнозування подій та генерацію сценаріїв. Застосування платформи трансферу знань дозволяє інтегрувати та надійно зберігати дані, аналізувати вплив ключових факторів та забезпечувати адаптивність до змінних умов. Наукова новизна дослідження полягає у розробці нових методів моделювання, використанні штучного інтелекту для аналізу даних та впровадженні інструментів інтеграції знань. Представлені результати є важливими для прийняття стратегічних рішень у державному управлінні, військовій справі та бізнесі в умовах невизначеності.

Ключові слова: автоматизовані рішення, трансфер знань, сценарії, штучний інтелект, моделювання, стратегічне планування, невизначеність.

Актуальність дослідження

У сучасних умовах невизначеність є ключовим фактором, що впливає на прийняття рішень. Це створює необхідність розробки сценаріїв можливих подій із використанням автоматизованих рішень. Актуальність таких підходів особливо зростає у

середньо- та довготерміновому прогнозуванні, державному управлінні, військовій сфері та бізнесі.

Автоматизація процесів генерації сценаріїв за допомогою платформ трансферу знань забезпечує не лише ефективність, але й адаптивність до змінних умов. Це дозволяє аналітикам зосередитися на глибокій інтерпретації даних та виявленні альтернативних можливостей розвитку подій.

Проблематика та постановка задачі

Головною метою розробки сценаріїв є створення моделей можливого розвитку подій із використанням автоматизованих рішень. Ці моделі не є передбаченнями, але вони враховують вплив ключових факторів, демонструючи різні сценарії майбутнього.

Процес генерації сценаріїв повинен базуватися на:

- Використанні платформ трансферу знань для інтеграції даних та знань аналітиків.
- Застосуванні індикаторів, які сигналізують про зміну умов або підтверджують певні події.
- Експертному оцінюванні факторів впливу.

Це забезпечує створення структурованих сценаріїв, що охоплюють широкий спектр можливих альтернатив.

Хід рішення

Відповідно до вимог і методів групової декомпозиції [1] та експертного оцінювання з підтримкою прийняття рішень [2, 3], сформовано чітку послідовність дій, яка дозволяє аналітичній групі систематично створювати реалістичні сценарії та автоматизовано генерувати їх у необхідній кількості для задоволення потреб замовника.

Запропонована в [4] методологія розробки сценаріїв базується на двоетапному процесі:

Етап 1: Формування прогнозованих подій

1. Визначення цілей проблемної ситуації.
2. Декомпозиція проблеми для виділення ключових сфер діяльності(за допомогою технології, реалізованої в системі «Консенсус-2» [5]).
3. Збір інформації з бази даних учасників процесу.
4. Побудова ієрархічної моделі впливів.

5. Рейтингування факторів впливу за допомогою СППР («Солон-3» [6]).

6. Формування набору альтернативних подій.

Етап 2: Генерування сценаріїв

1. Вибір цілей, релевантних для кожного гравця.

2. Оцінка впливу подій на досягнення цілей.

3. Декомпозиція кожного сценарію.

4. Створення деталізованої послідовності подій.

Усі етапи виконуються з використанням автоматизованих рішень на платформі трансферу знань, що забезпечує систематизацію та підвищення якості прогнозів.

Наукова новизна

1. Розроблено новий метод розробки сценаріїв на основі моделей проблемних ситуацій, що враховують причинно-наслідкові зв'язки між подіями.

2. Впроваджено інструментарій для отримання знань аналітиками з різних джерел інформації.

3. Запропоновано використання штучного інтелекту для автоматизації процесів аналізу та підвищення обізнаності аналітиків.

Висновки

Розробка сценаріїв із використанням автоматизованих рішень на платформах трансферу знань є важливим інструментом для стратегічного планування. Вона сприяє покращенню якості прийняття рішень у сферах з високим рівнем невизначеності, таких як державне управління, військова справа та бізнес.

Такий підхід може стати основою для інноваційного управління майбутніми викликами.

1. Тоценко В.Г. Методи і системи підтримки прийняття рішень. Алгоритмічний аспект / ІППІ НАНУ. – К.: Наукова думка, 2002. – 382 с.

2. Totsenko V.G. One Approach to the Decision Making Support in R&D Planning. Part 2. The Method of Goal Dynamic Estimating of Alternatives. Journal of Automation and Information Sciences. – 2001. – vol. 33, #4. – P. 82-90.

3. Циганок В.В. Удосконалення методу цільового динамічного оцінювання альтернатив та особливості його застосування. Реєстрація, зберігання і обробка даних. 2013, т.15, №1 – с. 90-99.

4. Циганок В.В., Роїк П.Д. Технологія генерування сценаріїв в умовах невизначеності / Реєстрація, зберігання і обробка даних: зб. наук.

праць за матеріалами Щорічної підсумкової наукової конференції 28 вересня 2020 року / НАН України. Інститут проблем реєстрації інформації. – К: ІПРІ НАН України, 2020. – С. 119-121.

5. Свідectво про реєстрацію авторського права на твір №75023. Комп'ютерна програма „Система розподіленого збору та обробки експертної інформації для систем підтримки прийняття рішень – «Консенсус-2»” / Циганок В.В., Роїк П.Д., Андрійчук О.В., Каденко С.В. // від 27/11/2017.

6. Свідectво про державну реєстрацію авторського права на твір №8669. Міністерство освіти і науки України державний департамент інтелектуальної власності. Комп'ютерна програма "Система підтримки прийняття рішень СОЛОН-3" (СППР СОЛОН-3) / Тоценко В.Г., Качанов П.Т., Циганок В.В.// зареєстровано 31.10.2003.

ГРАФОВА МОДЕЛЬ МЕНТАЛЬНОЇ ВІЙНИ

Анатолій Качинський¹, ORCID: [0000-0001-9642-7006],
akachynsky@gmail.com

Дмитро Ланде^{1,2}, ORCID: [0000-0003-3945-1178],
dwlande@gmail.com

¹КПІ ім. Ігоря Сікорського, ²ІПРІ НАН України, Київ, Україна

Метою цієї роботи є створення моделі «ментальних воєн». Для модернізації базової ієрархічної моделі використовується генеративний штучний інтелект. Запропонована методологія передбачає розширення моделі новими концептами, категоріями та зв'язками між ними, що генеруються за допомогою систем штучного інтелекту та великих мовних моделей. Застосовуючи кластеризацію на основі модулярності та ранжирування вузлів у мережевій моделі «ментальних воєн», вдосконалюється та розширюється базова ієрархічна модель, виявляють нові аспекти та поглиблюють розуміння змісту, цілей і наслідків інформаційно-психологічних ментальних воєн.

Ключові слова: Ментальні війни, графова модель, генеративний штучний інтелект, великі мовні моделі, семантична мережа, кластеризація

Вступ

У сучасному контексті реальної війни росії проти України дедалі важливішим стає врахування не лише традиційних методів

негативного впливу пропаганди, а й сучасних технологічних інструментів. Мета цієї роботи – розширити й удосконалити модель ментальних війн з урахуванням сучасного стану інформаційно-психологічного протистояння за допомогою генеративного штучного інтелекту (ШІ) та великих мовних моделей (LLMs).

Останніми роками генеративний штучний інтелект [1-3] справив значний вплив на когнітивні й інформаційні теорії, відкриваючи нові можливості для аналізу та моделювання складних систем. Використання ШІ не лише дозволяє вдосконалювати наявні моделі, а й виявляти нові аспекти, які раніше залишалися поза увагою дослідників [4].

Категорію "ментальна війна" ми розглядаємо як складну систему з ієрархічною структурою [5]. Нехай H це скінченна частково впорядкована множина з найбільшим елементом b . H є ієрархією, якщо виконуються наступні умови:

1. Існує розбиття множини H на підмножини L_k , $k = 1, 2, \dots, h$, де $L_1 = \{b\}$.
2. З того, що $x \in L_k$, випливає $x^- \in L_{k+1}$, $k = 1, 2, \dots, h-1$.
3. З того, що $x \in L_k$, випливає $x^+ \in L_{k-1}$, $k = 2, \dots, h$.

Тут x^- – це елемент, що йде після елемента x , а x^+ – це елемент, що передує елементу x у рівнях ієрархії.

При такому розгляді проблеми можна розкласти на простіші компоненти, після чого оцінити відносний ступінь взаємодії між елементами цієї ієрархічної структури. У складній системі "ментальної війни" можна виділити п'ять рівнів L_i , де $i=1,2,\dots,5$. Традиційний погляд на ментальну війну як ієрархічну структуру передбачає такі ролі для кожного рівня: на першому рівні (L_1 — цілі ментальної війни) розглядається один елемент — фокус, який розташований на вершині ієрархії (війна за трансформацію ідентичності); на другому рівні (L_2 — сили та засоби ментальної війни) відображаються економічні, політичні та соціальні сили, що впливають на результат (фінанси, література, мистецтво загалом, медіа, Інтернет тощо); третій рівень (L_3 — актори ментальної війни) складається з акторів, які маніпулюють цими силами (уряд, митці, меценати тощо); четвертий рівень (L_4 — цілі акторів) визначає завдання кожного актора (зміна сприйняття, цінностей, установок, стереотипів, традицій, архетипів національної свідомості); п'ятий рівень (L_5 — політики, що реалізуються акторами) описує можливі сценарії або результати,

до яких прагнуть актори через свої політики, зокрема перекодування цивілізаційної ідентичності держави, а також культурних цінностей суспільства та окремих осіб.

Ієрархічна декомпозиція дозволяє структурувати систему на підсистеми, кожна з яких відповідає за конкретні завдання. Математична модель допомагає формалізувати взаємодії між підсистемами за допомогою графів, матриць зв'язків та цільових функцій, які описують загальну ефективність системи.

Рівні L_1 та L_4 ієрархії включають певні набори цілей і підцілей. Нехай на рівні i є набір цілей $T_i = \{T_{i1}, T_{i2}, \dots, T_{jmi}\}$, де $j = \{1, 4\}$, а mi — кількість цілей на рівні i . Кожній цілі T_{ij} відповідає набір підцілей $F_{ij} = \{F_{ij1}, F_{ij2}, \dots, F_{ijn}\}$, де ijn — кількість підцілей для цілі T_{ij} . Таким чином, глобальною метою першого рівня T_{1i} є остаточне й незворотне розчинення української ідентичності в «загальноросійській» та відмова українців як політичної нації від претензій на державність за відсутності усвідомлення власної унікальності. У той же час цілі акторів T_{4i} , залишаючись підпорядкованими глобальній меті «ментальної війни» Росії проти України, можуть змінюватися залежно від епохи, переважно за формою і деякими характеристиками, залишаючись незмінними за своєю суттю.

Елементи та підцілі на кожному рівні виконують конкретні функції (інформаційні, економічні, організаційні тощо) і можуть бути пов'язані з іншими елементами та підцілями F_{ijk} на тому ж або інших рівнях, утворюючи мережу взаємозв'язків. Це описується набором функціональних зв'язків між елементами та підцілями:

$$Links = \{(F_{ijk}, F_{i'j'k'}) \mid F_{ijk} \text{ is linked to } F_{i'j'k'}\}.$$

Набір функціональних зв'язків між елементами, цілями та підцілями формує орієнтований граф, у якому вершинами є елементи, підцілі та цілі F_{ijk} , а ребрами — зв'язки між ними. Позначимо це як $G = (V, EG)$, де V — набір вершин, а $EG \subseteq V \times V$ — набір орієнтованих ребер, що описують функціональні залежності між елементами, підцілями та цілями. Для кожного елемента системи F_{ijk} визначається функція $f(F_{ijk})$, яка описує його роль у системі. Функції можуть бути представлені як набір залежностей або формул, що описують, як ці елементи взаємодіють один з одним у межах мережі. Зв'язки між елементами, цілями та підцілями на різних рівнях можна представити через матрицю зв'язків між рівнями. Нехай $A_{ii'}$ — це

матриця розміром $m_i \times m_i$, де кожен елемент $a_{ij}, a_{ij'}$ описує наявність або відсутність зв'язку між цілями та підцілями T_{ij} і $T_{ij'}$.

Для моделювання процесу досягнення головних цілей на кожному рівні можна визначити цільову функцію системи «ментальної війни» Φ , яка залежить від виконання цілей і підцілей на різних рівнях:

$$\Phi = \alpha f(T_{ij}) + \beta (F_{ijk}),$$

де α, β — вагові коефіцієнти у межах інтервалу зі значеннями в діапазоні $[0,1]$.

Перехід від ієрархічної моделі до мережевої

У контексті ієрархічної структури ментальних воєн як підцілі, так і концепти на відповідному рівні можуть визначатися експертами, включаючи віртуальних експертів [3], які допомагають створювати та вдосконалювати систему цілей та їхніх відповідних концептів. Якщо окремі концепти можуть одночасно належати до кількох різних рівнів ієрархії, то система зв'язків перетворюється на мережеву. Експерти, включно з віртуальними, визначають ці концепти на основі аналізу рівнів у первинній ієрархії моделі ментальної війни. Концепти можуть одночасно належати до декількох рівнів (включаючи цілі, визначені вище), що ускладнює зв'язки між рівнями моделі. Якщо концепт f_i належить до кількох рівнів одночасно, це вводить нові зв'язки між рівнями. Таким чином, структура системи змінюється з чистої ієрархії на мережеву модель. Якщо поняття f_i належить рівням T_a та T_b , то між цими рівнями існує зв'язок:

$$E = \{(T_a, T_b) \mid f_i \in F, f_i \in T_a \cap T_b\}.$$

Ця формалізація відповідає послідовності дій, що передбачає можливість повторення процесу під наглядом експерта до досягнення остаточного розуміння предметної області:

1. Формування початкової схеми із відображенням базових зв'язків між поняттями.
2. Створення промптів для LLM з метою генерації нових понять і зв'язків.
3. Інтеграція нових зв'язків у початкову схему.
4. Обробка лінгвістичних даних – оцінювання нових та існуючих зв'язків, а також ранжування вузлів.

5. Аналіз даних і візуалізація шляхом завантаження даних у систему аналізу графів.
6. Формування та уточнення кластерів, визначення їхніх назв за допомогою LLM.
7. Фінальна перевірка та валідація розширеної моделі.

Аналіз і візуалізація

Об'єднані дані завантажуються в середовище програми Gephi для кластеризації на основі класів модулярності (Рис. 1). У Gephi можна застосовувати різні типи модулярності, зокрема, модель Поттса [6], яка враховує так званий параметр роздільної здатності. Згідно з цією моделлю, функція якості $H(G,P)$, або скорочено $H(P)$, для розбиття P на модулі (кластери) графа G записується так:

$$H(P) = -\sum_C e_C - \gamma n_C^2,$$

де кожен кластер $C \in P$ складається з ребер e_C і вузлів n_C , а γ є параметром роздільної здатності, який суттєво впливає на поділ графа на кластери.

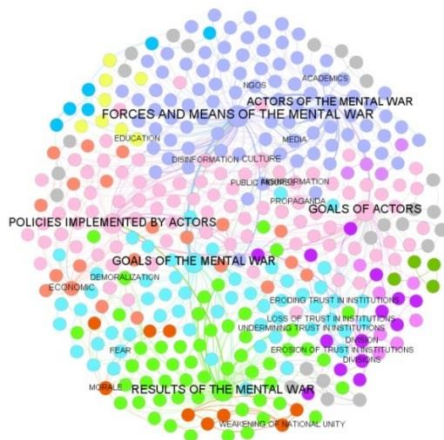


Рисунок 1 – Мережа, що відповідає розширеній моделі ментальних воєн

Процедура визначення класів модулярності виконується за таким алгоритмом, який охоплює наступні етапи:

1. Початковий розподіл вузлів на початкові кластери.

2. Визначення модулярності для поточного розподілу вузлів.

3. З'єднання груп вузлів для перевірки, чи покращується значення модулярності.

4. Повторення процесу об'єднання кластерів та оцінки модулярності, поки не буде досягнуто її максимальне значення.

Висновки

Розширена модель ментальних війн демонструє комплексну мережу взаємозв'язків між основними та додатковими концептами. Нові концепти, виявлені за допомогою LLM, дозволили глибше зрозуміти категорію "ментальних війн". Кластеризація та ранжирування вузлів дозволяють виокремити ключові зв'язки та концепти, що впливають на розуміння аспектів ментальних війн. Семантичну мережу було розділено на декілька кластерів, кожен з яких представляє окремий аспект ментальних війн. Алгоритм кластеризації на основі класів модулярності дозволив визначити значущі зв'язки між вузлами та встановити ключові компоненти мережі. Ранжирування вузлів допомогло ідентифікувати найбільш впливові концепти в моделі. Це дозволило визначити, які елементи є найважливішими для подальшого вивчення та вдосконалення.

1. Dmytro Lande, Leonard Strashnoy. *GPT Semantic Networking: A Dream of the Semantic Web - The Time is Now*. Kyiv: Engineering, 2023. ISBN 978-966-2344-94-3.

2. Srinivasan Ramanujam. *The LLM Revolution: Transforming Industries with Large Language Models*. Kindle Edition, 2024. <https://www.amazon.com/LLM-Revolution-Transforming-Industries-Language-ebook/dp/B0D7VS5BNK>

3. Dmytro V. Lande, Leonard Strashnoy. *Causality Network Formation with ChatGPT*. SSRN Electronic Journal. (May 30, 2023). – 16 p. DOI: 10.2139/ssrn.4464477

4. Dmytro Lande. OSINT in Cybersecurity: Educational Manual. – Kyiv: LLC "Engineering", 2024. – 522 p.

5. Anatolii Kachynskiy. The structural-functional model of the system for ensuring informational and informational-psychological security. Reports of the National Academy of Sciences of Ukraine, 2023. №1. pp. 16-23. (ukr. lang)

6. F. Y. Wu. *The Potts model*. Rev. Mod. Phys. 54, 235 – Published 1 January 1982

УНІВЕРСАЛЬНИЙ МЕТОД ВИЗНАЧЕННЯ РІВНЯ УЗГОДЖЕНОСТІ ЕКСПЕРТНИХ ПОПАРНИХ ПОРІВНЯНЬ, ДОСТАТНЬОГО ДЛЯ ЇХ АГРЕГАЦІЇ

Павло Роїк¹, Андрій Оленко²,
Сергій Каденко¹, Олег Андрійчук^{1,3}

¹ Інститут проблем реєстрації інформації Національної академії
наук України, Київ, Україна

² La Trobe University, Victoria 3086, Melbourne, Australia

³ Київський національний університет імені Тараса Шевченка,
Київ, Україна

*paulroyik@gmail.com; a.olenko@latrobe.edu.au;
seriga2009@gmail.com; andriichuk@ipri.kiev.ua*

Розроблено інноваційний метод визначення порогу узгодженості експертних попарних порівнянь, який дозволяє обчислювати його значення залежно від необхідного рівня достовірності експертизи. Метод базується на імітаційному моделюванні та генетичному алгоритмі пошуку, забезпечує об'єктивну оцінку узгодженості, підвищення якості експертиз та ефективне використання експертних ресурсів у системах підтримки прийняття рішень.

Ключові слова: узгодженість експертних оцінок, достатній для агрегації рівень узгодженості, моделювання оцінок, генетичний алгоритм

Актуальність дослідження

В сучасних системах підтримки прийняття рішень (СППР) експертне оцінювання відіграє критично важливу роль. Ключовою проблемою залишається забезпечення достовірності експертних оцінок, яка безпосередньо залежить від узгодженості суджень експертів. Традиційні підходи визначення узгодженості мають суттєві обмеження, оскільки не пов'язують рівень узгодженості з вимогами до достовірності результатів експертизи. Недостатній рівень узгодженості може призводити до недостовірності результатів, а отже, до прийняття неправильних рішень.

Одним із етапів аналізу рівня узгодженості є визначення порогу узгодженості експертних оцінок, нижче якого немає сенсу у подальшому аналізі експертних оцінок.

Цей етап є надзвичайно важливим через кілька ключових факторів:

1. Підвищення достовірності моделей: У СППР результати експертних оцінок є основою для побудови моделей слабо структурованих предметних галузей. Низький рівень узгодженості може спотворювати результати агрегації, що безпосередньо впливає на якість рекомендацій для осіб, що приймають рішення.

2. Запобігання інформаційним конфліктам: У процесі групових експертиз часто виникають протиріччя між оцінками експертів. Визначення порогу узгодженості дозволяє ідентифікувати випадки, коли агрегація недоцільна, і вимагає додаткового уточнення оцінок, що знижує ризик прийняття хибних рішень.

3. Економія ресурсів: Визначення порогу узгодженості дозволяє мінімізувати кількість необхідних ітерацій для досягнення прийняттого рівня узгодженості, тим самим оптимізуючи процес використання експертних ресурсів і часу.

4. Забезпечення об'єктивності: Встановлення чітких критеріїв узгодженості сприяє стандартизації процесу експертизи, що знижує вплив суб'єктивних чинників і підвищує довіру до отриманих результатів.

Методологія дослідження

Запропонований підхід [1] базується на концепції імітаційного моделювання експертних попарних порівнянь з використанням генетичного алгоритму [2]. Основні етапи методології включають:

1. Визначення гіпотетичних еталонних ваг альтернатив.
2. Моделювання матриць попарних порівнянь (МПП) з введенням контрольованої похибки.
3. Цілеспрямований пошук найменш узгоджених варіантів МПП за допомогою генетичного алгоритму.
4. Встановлення залежності між похибкою оцінювання та рівнем узгодженості.

Було запропоновано метод, який в залежності від заданого бажаного рівня достовірності результатів експертизи обчислював значення порогу узгодженості.

Цей метод ґрунтується на збуренні ідеально узгоджених МПП та знаходженні найбільш неузгодженого випадку. Далі, суть методу подається у спрощеному вигляді.

Задаються довільні позитивні значення ваг альтернатив $w_i, i = \overline{1..n}$. Число альтернатив n при експертному оцінюванні не перевищує 7 ± 2 . Проводиться нормування цих ваг до одиниці $w_i = \frac{w_i}{\sum_{j=1}^n w_j}$, щоб виконувалась умова $\sum_{i=1}^n w_i = 1$. Надалі ці ваги вважаються еталоном. За цими еталонними вагами формально будується мультиплікативна ідеально узгоджена МПП A виходячи зі співвідношення $a_{ij} = \frac{w_i}{w_j}$, де a_{ij} – елемент матриці A .

Далі, проводиться «зашумлення» матриці A , таким чином, що кожен елемент матриці A , крім діагональних, може бути змінений за наступним законом: $a'_{ij} = a_{ij} \pm a_{ij} \cdot \delta / 100\%$, де $0 < \delta < 100$. Величина δ - наперед задана величина, що характеризує максимальне відносне відхилення елементів матриці A у відсотках від еталонних значень. Тим самим, здійснюється імітаційне моделювання похибок при експертному оцінюванні. Варіюючи δ , знаходиться індекс узгодженості для найбільш неузгодженого випадку МПП. Це й буде поріг узгодженості для певного значення δ .

Ключові результати

Дослідження дозволило:

- Експериментально встановити залежність максимального відхилення агрегованих оцінок від еталонних значень;
- Визначити кореляцію між відносною похибкою експертного оцінювання та узгодженістю матриць порівнянь;
- Розробити універсальний підхід визначення порогу узгодженості.

Практичне значення

Запропонована методологія надає кількісний інструментарій для:

- Об'єктивного встановлення порогу узгодженості експертних оцінок;
- Підвищення достовірності результатів експертизи;
- Ефективного використання експертних ресурсів.

Експериментальне дослідження проведено для нового запропонованого індексу узгодженості [3] для діапазону відносних похибок експертного оцінювання від 0 до 50%, що репрезентативно відображає реальні експертні практики.

Окрім того, метод має потенціал універсального застосування для різних інших індексів узгодженості, зокрема:

- Коефіцієнта неузгодженості методу аналізу ієрархій Т. Сааті [3];
- Індекс узгодженості подвійної ентропії [4];
- Спектральних методів визначення узгодженості [5].

Висновки

Представлений метод є інноваційним кроком у вирішенні проблеми забезпечення якості експертних оцінок, надаючи кількісний інструментарій для об'єктивного визначення порогу узгодженості при парних порівняннях.

1. VitaliyTsyganok, AndriyOlenko, PavloRoik, OksanaVlasenko, Determining adequate consistency levels for aggregation of expert estimates. *arXiv preprint. Methodology (stat.ME)* arXiv:2410.03012v1 [stat.ME] 3 Oct 2024 9p. <https://doi.org/10.48550/arXiv.2410.03012>
2. Tsyganok V.V. Investigation of the aggregation effectiveness of expert estimates obtained by the pairwise comparison method. *Mathematical and Computer Modelling.*– August 2010. – v.52, №3-4. – P.538-544.DOI: 10.1016/j.mcm.2010.03.052
3. Saaty T.L. & Kearns K.P. Analytical Planning the Organization of Systems, *Pergamon Press*, 1985.
4. Olenko Andriy and Tsyganok Vitaliy, Double Entropy Inter-Rater Agreement Indices. *Applied Psychological Measurement.* 2016. Vol. 40(1). P. 37–55. DOI: 10.1177/0146621615592718
5. Totsenko V.G. Spectral Method for Determination of Consistency of Expert Estimate Sets. *Engineering Simulation.* – 2000. – 17. – P.715-727.

LIKELIHOOD ESTIMATION - BASED POST - FILTERING FOR GENERALIZED SIDELobe CANCELLER BEAMFORMER

Quan Trong The

Faculty of Information Security

Posts and Telecommunications Institute of Technology

Hanoi, Vietnam

theqt@ptit.edu.vn, [0000-0002-2456-9598]

Speech enhancement technique is widely integrated into numerous acoustic devices with the aim of improve the speech quality, the obtained perceptual metric listener, the overall

speech intelligibility of the captured signal, which often corrupted by complex background noise, unwanted interference, competing talker. Recently, microphone array (MA) beamforming widely applied into various types of speech applications, such as voice - controlled equipments, hearing aids, teleconference system, smart-home, hands-free mobile phone for suppressing surrounding noise while extracting the original speech component. Generalized Sidelobe Canceller (GSC) beamformer has attracted a lot of attentions due to its high performance in the human - machine interactive system for achieving better noise reduction, speech enhancement. In many acoustic applications, such as hearing aids, voice - controlled vehicle, cochlear implant devices, speech recognition, GSC beamformer is very sensitive with the precise estimation of preferred steering vector, the microphone sensitivities, the environmental homogeneity of recording scenario. Because of the inaccurate calculation of these factors, GSC beamformer's performance often degraded, that still exits unwanted noise level at the output of beamformer. In this contribution, the author proposed an effective post - Filtering, which based on the likelihood estimation of speech variance according to the changed recording situation. The demonstrated experiment with dual - microphone system (DMA2) has confirmed the effectiveness of this approach in reducing the noise level to 12.4 dB in comparison with traditional GSC beamformer and increasing the signal-to-noise (SNR) ratio from 8.1 to 9.2 dB. The promising result has shown the significant improvement of GSC beamformer's evaluation by applying the post - Filtering in adverse recording environment.

Index Terms - microphone array beamforming, noise reduction, speech enhancement, generalized sidelobe canceller, the signal-to-noise ratio, likelihood estimation, speech variance, post - Filtering

BIOTHREAT EARLY ASSIST AND RESPONSE COMMAND SYSTEM (BEAR-CS): A STATE-OF-THE- ART ANALYSIS AND COMPARATIVE STUDY WITH EXISTING BIO SURVEILLANCE AND INCIDENT MANAGEMENT SYSTEMS

**Rexhep Mustafovski¹, Aleksandar Risteski² and Tomislav
Shuminoski³**

*^{1,2,3}University "St. Cyril and Methodius" – Skopje, Faculty of
Electrical Engineering and Information Technologies, street. Ruder
Bosković, 1000 Skopje, Republic of North Macedonia*

Biological threats, such as viruses and hazardous agents, present significant risks to both military operations and civilian safety. The Biothreat Early Assist and Response Command System (BEAR-CS) is an innovative, integrated command system aimed at proactively identifying biothreats in operational areas prior to the deployment of military personnel. This system employs drones equipped with advanced sensors to monitor regions and send real-time data to Tactical Operating Centres (TOCs), which then communicate this information to command posts in the field. Soldiers are provided with enhanced protective technologies, including digital radios, triage tools, poison detection devices, and decontamination kits, to ensure their safety and the success of their missions. This paper compares BEAR-CS with existing systems for incident management, bio surveillance, and threat detection, as outlined in three key studies. These studies showcase the capabilities of the Incident Command System (ICS) for Emergency Medical Services, the BioWatch and BD21 bio surveillance systems, and the Countering Weapons of Mass Destruction (CWMD) framework. A thorough comparative analysis highlights the unique benefits of BEAR-CS, such as real-time detection, pre-incident threat assessment, and seamless integration with soldier systems. Additionally, this paper explores the limitations of current systems, including delayed detection times, restricted deployment flexibility, and a lack of pre-emptive measures. The findings underscore how BEAR-CS addresses these shortcomings and suggest future improvements, such as the integration of artificial intelligence and expanded sensor networks. This cutting-edge analysis

emphasizes the potential of BEAR-CS to transform biothreat detection and improve military effectiveness in complex operational settings.

Keywords: Biothreat detection, BEAR-CS, biosurveillance, drone-based monitoring, Tactical Operating Centers (TOCs), Incident Command System (ICS), BioWatch, BD21, Countering Weapons of Mass Destruction (CWMD), military operations

СТРУКТУРНА СКЛАДНІСТЬ ПОРОГОВИХ ЯВИЩ У СКЛАДНИХ МЕРЕЖАХ

А.О. Снарський^{1,2}, І.В.Дідур²

¹*Institute for Information Recording NAS of Ukraine,*

²*National Technical University of Ukraine “Igor Sikorsky Kyiv Polytechnic Institute”*

asnarskii@gmail.com, ivan@didur.io

Вступ

Одним із ключових викликів сучасної науки є розробка універсальних методів аналізу даних, які можуть бути представлені у вигляді графів або мереж. Такі підходи є основою для моделювання складних систем у фізиці, біології та інформатиці, що охоплюють фазові переходи, динаміку біологічних мереж та функціонування нейронних мереж [1-5]. У цій роботі ми пропонуємо підхід до кількісного аналізу мереж шляхом введення скалярної метрики. Ця метрика дозволяє оцінити складність мережі, що в свою чергу надає можливість відранжувати будь-які дані, які можна представити у вигляді мережі, забезпечуючи універсальний інструмент для оцінки структурної складності систем.

Поняття структурної складності складних мереж було введено у як аналог структурної складності зображень [2]. Структурна складність є важливою характеристикою, яка описує як загальну організацію мережі, так і локальні властивості її окремих вузлів. Ця характеристика походить з методів ренормалізації у їхньому геометричному застосуванні, а також з методів Detrended Fluctuation Analysis [3]. Завдяки цьому підходу можна аналізувати, як змінюється структура складних мереж на різних масштабах, зокрема через введення поняття масштабу для складних мереж.

Попередні дослідження [1] показали, що структурна складність дозволяє визначати в тому числі критичні значення концентрації для фазових переходів у зростаючих графах Ердоша-Реньї, а також аналізувати поведінку графа видимості, побудованого для RR-інтервалів серцевих ритмів при різних аритміях. У цій роботі пропонується новий підхід до оцінки структурної складності. На відміну від [1], де підстилаючий граф визначався як повнозв'язний, у цьому дослідженні обрані графи з неповнозв'язною структурою, наприклад, у вигляді квадратних або трикутних решіток.

У роботі продемонстровано приклади аналізу перколяційних мереж на квадратній, трикутній та шестикутній решітках. Введена методика дозволяє виявляти критичні значення параметрів для фазових переходів у таких системах, забезпечуючи універсальний підхід до вивчення складних мереж.

Результати дослідження мають широкий спектр практичних застосувань, включаючи класифікацію складних мереж, ідентифікацію критичних станів у фізичних та біологічних системах, кореляцію метрик різних нейронних мереж з їх структурною складністю, а також оптимізацію обчислювальних процесів для великих мереж із використанням методів паралельного програмування [5, 6].

Основні положення

Розрахунок складності в [2] базується на чотирьох ключових принципах. По-перше, методика повинна агрегувати інформацію з різних масштабів графа, забезпечуючи мультишаровий аналіз. По-друге, формула розрахунку є аналітично визначеною, що мінімізує необхідність суб'єктивних рішень під час обчислень. Третій принцип полягає в стабільності метрики: вона повинна залишатися стійкою до невеликих змін у структурі графа, що забезпечує її надійність для практичних застосувань. Нарешті четвертий нормує величину складності таким чином, що крайніми, нульовими значеннями є порожні та повні графіки.

У цьому дослідженні ми пропонуємо методику оцінки складності графів на основі введення нової скалярної метрики, яка узагальнює запропоновану в [2] та оптимізована для великих структур за допомогою CUDA [8-9]. Наш підхід базується на узагальненому алгоритмі обчислення складності, який встановлює нульові значення складності для повного графа та так званого базового графа, який обирається окремо для різних задач.

Алгоритм розрахунку

Алгоритм розрахунку складності починається з підготовки даних, де вхідними даними є граф із заданою структурою, наприклад, квадратна, трикутна, шестикутна решітка, або довільна (програмне забезпечення дозволяє користувачу намалювати граф), а також параметр ймовірності зв'язку p .

На першому етапі для кожної вершини графа обчислюється локальна складність, що враховує кількість зв'язків у вершині, кількість заповнених зв'язків та їх максимальну можливу кількість на різних масштабах.

Наступним кроком є агрегування локальних значень складності для всього графа. Це досягається обчисленням середнього значення локальних складностей для всіх вершин графа.

Останній етап полягає в оптимізації розрахунків, яка виконується з використанням CUDA [8] для ефективного паралельного виконання на графічних процесорах, що суттєво знижує час обчислень при роботі з великими графами.

Результати

Тут, як приклад, розглянуто поведінку структурної складності на трьох нескінченних ґратках — квадратній, трикутній та шестикутній — із заданою ймовірністю існування (відповідно розриву) зв'язку (задача перколяції, рис. 1 [7]). У кожній із цих ґраток частина зв'язків зруйнована, а ймовірність наявності зв'язку позначена як p . Нульовою складністю, у запропонованому визначенні, є ґратка при $p = 0$, та при $p = 1$.

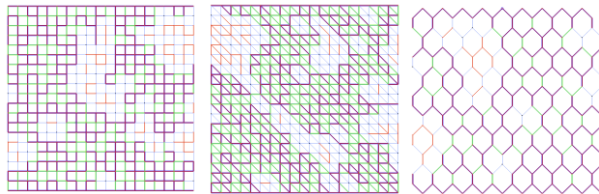


Рис. 1. Зображення решіток на відповідних порогах перколяції.

Для квадратної $p_c = 0.5$, для трикутної $p_c = 0.347296$, для шестикутної $p_c = 0.652703$.

Найбільш складною структурою в перколяційній системі є так званий нескінченний кластер [7]. Він має, зокрема, фрактальні

властивості. Концентрація, при якій спостерігається нескінченний кластер, називається порогом протікання - p_c . Отже, запропонована характеристика структурної складності повинна мати в цій точці концентрації особливе (характеристичне) значення.

На рис. 2 показано, як складність змінюється залежно від параметра ймовірності зв'язку (концентрації) для кожного типу решітки. Точками показані значення складності для заданої концентрації при різних реалізаціях випадкової структури (розподілу розірваних зв'язків), а суцільна лінія — апроксимація, побудована за їх середніми значеннями.

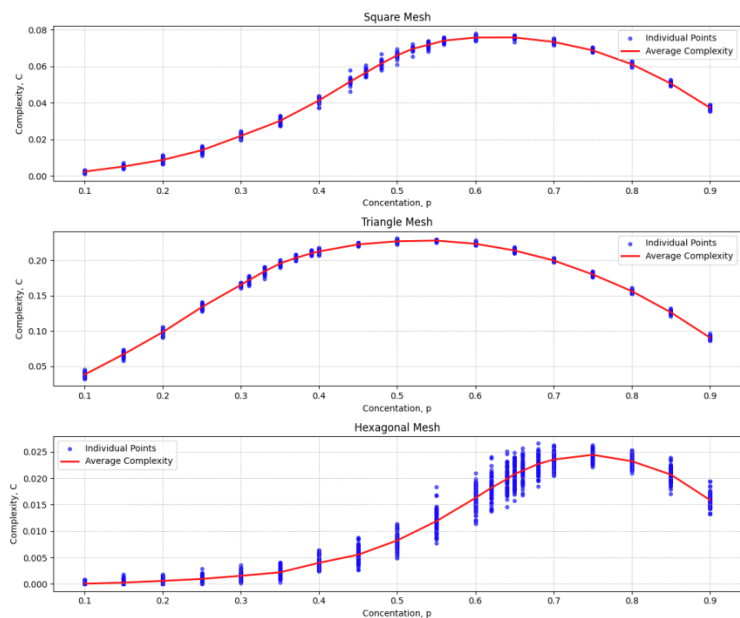


Рис. 2. Зображення складності в залежності від ймовірності зв'язку

Результати продемонстрували, що похідна складності збігається (вказує) на значення порогу протікання: для квадратної $p_c \sim 0.5$, для трикутної $p_c \sim 0.35$, для шестикутної $p_c \sim 0.65$. Це збігається з літературними даними [1, 2] та підтверджує валідність метрики.

Висновки

У цій роботі було розроблено нову скалярну метрику для оцінки структурної складності графів із довільною топологією, яка дозволяє аналізувати фізичні, біологічні та інформаційні системи. На основі проведеного аналізу підтверджено, що метрика є валідною та стабільною для трьох типів решіток (квадратної, трикутної, шестикутної). Таким чином, для введеної метрики значущою характеристикою є значення її похідної.

Цитована література

1. Snarskii A. O. Transitions in Growing Networks Using a Structural Complexity Approach. *Physical Review E*. 2024. Vol. 110, no. 054309.
2. Bagrov A.A., Iakovlev I.A., Iliasov A.A., Katsnelson M.I., Mazurenko V.V. Multiscale Structural Complexity of Natural Patterns. *PNAS*. 2020. Vol. 117, no. 48. P. 30241–30251.
3. Peng C.-K., Buldyrev S.V., Havlin S., Simons M., Stanley H.E., Goldberger A.L. Mosaic Organization of DNA Nucleotides. *Physical Review E*. 1994. Vol. 49. P. 1685–1689.
4. Cabrera D., Arnold L., Cabrera L. The Simple Rules of Complex Networks. *Journal of Applied Systems Thinking*. 2020. Vol. 20, no. 5. P. 1–9.
5. Li A., Pan Y. Structural Information and Dynamical Complexity of Networks. *IEEE Transactions on Information Theory*. 2016. Vol. 62, no. 4. P. 2455–2469.
6. Málthe-Sørenssen A. Percolation Theory Using Python. *Lecture Notes in Physics*. Vol. 1029. Springer Nature Switzerland AG. 2024. ISBN: 978-3-031-59899-9. DOI: 10.1007/978-3-031-59900-2.
7. Stauffer D., Aharony A. *Introduction to Percolation Theory*. Revised edition. Taylor & Francis, 1992. 181 p. ISBN: 0748400273.
8. Quer S. A Parallel Many-core CUDA-based Graph Labeling Computation. SCITEPRESS. 2020. DOI: 10.5220/0009780205970605 In *Proceedings of the 15th International Conference on Software Technologies (ICSOFT 2020)*, pages 597-605 ISBN: 978-989-758-443-5.
9. Milišić H., Ahmić D., Sinanović H., Sarić E., Asotić A., Huseinović A. Parallelization Challenges of BFS Traversal on Dense Graphs Using the CUDA Platform. 2016 XI International Symposium on Telecommunications (BIHTEL). 2016. P. 24-26. DOI: 10.1109/BIHTEL.2016.7775712.

SURVIVAL MODELS AS COPULAS FOR GREEN RISKS PREDICTION

Kuznietsova Nataliia^{1,2}, ORCID: 0000-0002-1662-1974,
Kvashuk Illia¹, ORCID: 0009-0006-5585-3045,
Bidyuk Petro¹, ORCID: 0000-0002-7421-3565

¹ *National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute", 37, ave. Beresteyskyi, 03056, Kyiv, Ukraine*
² *Claude Bernard Lyon 1 University, 43 boulevard du 11 Novembre 1918, 69622 Villeurbanne cedex, Lyon, France*
natalia-kpi@ukr.net, illiakvashuk@gmail.com, pbidyuke_00@ukr.net

This work provides a comprehensive overview of the current state of copulas and proposes to use survival models for newly introduced families. While the copula's appearance has seen widespread application in various fields, including environmental sciences (particularly hydrology), economics, and as enhancements to existing models to improve performance. The extensive range of potential copulas, which can be systematically expanded through rigorously defined procedures such as those for the Archimedean class, underscores their versatility. However, many existing approaches emphasize forms of copulas that focus on past values, providing insights into the likelihood of events occurring before the value of interest. This paper introduces an alternative approach aimed at assessing the likelihood of events occurring in the future. This forward-looking perspective leverages the strengths of copulas more effectively, particularly in risk-sensitive fields such as environmental management.

Keywords: Copulas; Survival model; Archimedean class, green risks.

Introduction

Copulas are mathematical functions that bind marginal distributions to create a joint distribution, capturing dependencies between random variables. Its form is not dependent on the marginals due to the fact that arguments are values of CDFs and so allow for any underlying distribution. As per the definition and properties are given in [1], for any multivariable distribution, it is possible to write down

copulas that will catch the dependencies, in theory, and if proper copula is used, semi-perfectly. Independence of marginals and dependency form allows for a modular approach where individual and connection function are different, however, there are caveats that need to be taken into account [2].

There are several classes:

Archimedean copulas with special property [1][3]:

$$C(u_1, u_2, \dots, u_d) = \phi^{-1}(\phi(u_1) + \phi(u_2) + \dots + \phi(u_d)),$$

where ϕ is a generator function with special properties.

Vine copulas allow for the decomposition of dependencies on pair-wise form:

$$C(u_1, u_2, \dots, u_d) = \prod_{k=1}^{d-1} \prod_{i=1}^{d-k} c_{i,i+k|S}(u_i, u_{i+k}|u_S),$$

where $c_{i,i+k}$ is a copula that binds arguments u_i, u_{i+k} together with the impact of other parts of the model taken into account.

This approach is beneficial during scaling [4] when relationships between two variables have been modelled in advance.

Survival function as a copula

New copula can be constructed by directly following their definition or for some classes there is a defined procedure that streamlines the process. For Archimedean copula, it is possible to obtain a new copula by defining a special generator function, ϕ . Due to the weaker conditions than in a definition, it is possible to have a more robust approach for practical tasks. If some assumptions about the system's behaviour are available, it becomes feasible to generate entirely new copulas tailored to the problem at hand.

The approach with generation was used in the given work [3]. The idea was to use the half-logistic regression in order to create a model with only a positive range domain. Additionally, a new parameter was introduced to scale the model and improve its flexibility. The resulting copula is expressed as:

$$C_L(u, v; \theta) = \frac{1}{\theta} \ln \left[\frac{e^{\theta(1+u+v)} - e^{\theta u} - e^{\theta v} + e^{\theta}}{e^{\theta(1+u)} + e^{\theta(1+v)} - e^{\theta(u+v)} - 1} \right]$$

The newly obtained copula has shown its advantage on the hydrology dataset. However, little to no attention is given to the survival form.

The survival form of copula is derived by using survival function probability instead of cumulative. The formula is as follows:

$$\underline{C}(u_1, u_2, \dots, u_d) = u_1 + u_2 + \dots + u_d - d + C(1 - u_1, 1 - u_2, \dots, 1 - u_d)$$

In practical applications, this form is often more relevant for assessing the probabilities of extreme future events. For instance, in hydrology, one might be interested in predicting the likelihood of a variable (e.g., water flow) exceeding a certain threshold under specific conditions, such as high temperatures. Applying this transformation to the newly proposed copula yields:

$$\underline{C}(u, v; \theta) = \frac{\theta(u+v-1) + \theta + \ln\left(\frac{e^{2\theta} - e^{\theta u} - e^{\theta v} + e^{\theta(u+v)}}{-e^{2\theta} + e^{\theta(u+2)} - e^{\theta(u+v)} + e^{\theta(v+2)}}\right)}{\theta}$$

This formula is particularly useful for modelling the probability of extreme hydrological events, such as floods as it cumulative based on form has shown remarkably good results in overall modelling.

In our research, we decided to evaluate green risks as the probability in time which could be presented as a survival function. Thus, we also can determine our risk's probability as it was proposed in equation (1). It gives us the possibility to evaluate green risk as an environmental risk in the form of a copula as well as present it as a probability in the time interval via the survival model. It means that we will evaluate our green investment project as a function of some environmental and financial parameters that can change and vary during the time of consideration.

Conclusions

Overall, the advancements in the field of copulas are promising, both in terms of theoretical developments and practical applications. Existing tools have enabled researchers to generate copulas for a wide variety of scenarios, and real-world case studies have demonstrated impressive results. Fields such as environmental science and economics have particularly benefited from the implementation of copulas, where they have been instrumental in modelling complex dependencies and improving predictive accuracy.

The proposed copulas hold significant value, facilitating decision-making and enhancing risk management processes. By refining these approaches further, it is possible to make copulas more applicable to real-world scenarios while reducing the technical expertise required by users. Additionally, employing the survival form of copulas simplifies computations, allowing for more straightforward risk calculations and making these models more accessible for practical use.

[1] A. Sklar, "Functions de Répartition à n Dimensions et Leurs Marges", Publications de l'Institut de Statistique de l'Université de Paris, 8 (1959): 229–231.

[2] F. Tootoonchi, M. Sadegh, J. O. Haerter, O. Rätty, T. Grabs, C. Teutschbein, "Copulas for hydroclimatic analysis: A practice-oriented overview", *WIREs Water*, 9.2 (2022): e1579. doi:10.1002/wat2.1579.

[3] A. A. Alzaid, W. M. Alhadlaq, "A New Family of Archimedean Copulas: The Half-Logistic Family of Copulas", Department of Statistics and Operations Research, College of Science, King Saud University, Riyadh, Saudi Arabia. doi:10.3390/math12010101.

[4] P. Nagler, B. Pfaff, T. Brechmann, "Sparse vine copula regression: Handling high-dimensional dependence in statistical modelling", 2021. doi:10.1002/asmb.2043.

IDENTIFICATION OF PARAMETERS OF THE MATHEMATICAL MODEL OF A DYNAMIC OBJECT OF CRITICAL INFRASTRUCTURE

Kuchеров Dmytro, *orcid.org/0000-0002-4334-4175*, Kyiv, Ukraine,
d_kuchеров@ukr.net

National Aviation University

Fang Xiaopeng, *orcid.org/0009-0003-4043-0072*, Hangzhou,
Zhejiang, China, *fxproc@163.com*

Zhejiang University of Technology

Wang Zhiqiang, *orcid.org/0009-0003-7641-6176*, Hangzhou,
Zhejiang, China, *2066496033@qq.com*

Zhejiang University of Technology

Xu Zhongli, *orcid.org/0009-0000-6619-8468*, Shaoxing, Zhejiang,
China, *1422696998@qq.com*

Tongchuang Engineering Design Co., Ltd

Fu Xiaojian, *orcid.org/0009-0005-7393-2081*, Shaoxing, Zhejiang,
China, *271149355@qq.com*

Tongchuang Engineering Design Co., Ltd

The problem of determining the parameters of the mathematical model of a dynamic object by the method of identification by the acceleration curve is considered. The basis of the proposed identification method is the numerical integration of the plane of the error function. The report discusses the problems of implementing the algorithm for calculating the parameters of a dynamic object.

Keywords: dynamic object, numerical integration, identification of parameters

Introduction

When constructing automatic control systems for critical infrastructure objects, it is desirable to have as accurate a model of the controlled object as possible, which is not always possible. The structure and parameters of the mathematical model of a dynamic object are usually unknown. In this case, the designer faces a dilemma: to synthesize a system based on learning methods and adaptation to existing operating conditions or to perform identification, during which the structure and parameters of the dynamic object should be clarified. In the latter case, the procedure for determining the structure and parameters is somewhat simpler.

Features of method implementation

Identification of the structure and parameters of dynamic objects is carried out based on measuring the values of the transition function or the so-called acceleration curve obtained when a step-single function is applied to the object input. The estimated parameters are the delay time of the object's reaction, the gain coefficient, and the object inertia if the object is described by a first-order differential equation. If there are inflection points in the transition function, a conclusion is made about the order of the differential equation higher than the first order. In practical calculations, it is assumed that the order of the characteristic equation of the object is not higher than the third.

In some cases of designing automatic control systems, it is advisable to use a transfer function instead of a differential equation, which can completely replace it. Then, a suitable approach for determining the structure and parameters of a transfer function of the form

$$W_M(s) = \frac{1+b_1s+b_2s^2+\dots+b_ms^m}{1+a_1s+a_2s^2+\dots+a_ns^n} \quad (1)$$

the area method can be used, which allows one to find unknown coefficients of the transfer function (1) using the normalized acceleration curve $h(t)$ of the control object. The condition for the physical feasibility of (1) is the inequality $n \geq m$.

The method is based on the expansion in a Taylor series of the function inverse to (1) in the vicinity of the point $\xi = 0$ in the form

$$\frac{1}{W_M(s)} = \frac{1+a_1s+a_2s^2+\dots+a_ns^n}{1+b_1s+b_2s^2+\dots+b_ms^m} = 1 + S_1s + S_2s^2 + \dots + S_k s^k + \dots \quad (2)$$

At the same time $W_M^{-1}(0) = S_0 = 1$.

The coefficients S_k are found by subtracting the right-hand side of (2) from both sides of the equation, reducing to a common denominator, and equating the coefficients of s with the same powers. Performing this operation allows us to obtain the following expression

$$a_k = b_k + S_k + \sum_{i=1}^{k-1} b_i S_{k-i}. \quad (3)$$

Expression (3) contains $(n + m)$ unknown values, and therefore, to find the unknown coefficients a_k and b_k , it is necessary to compose independent equations for the variable S_k from (2), (3) to these coefficients, $k = 1, 2, 3, \dots$

To determine S_k , we consider the deviation function $\varepsilon(t)$ of the racing characteristic $h(t)$ from the input signal $1(t)$, as in (2) in Laplace images, i.e. $L\{\varepsilon(t)\}$ [1, p. 8.4]

$$L\{\varepsilon(t)\} = L\{1(t) - h(t)\} = L\{1(t)\} - L\{h(t)\} = \frac{1}{s} - W_m(s) \frac{1}{s} = \frac{1 - W_m(s)}{s} = E(s). \quad (4)$$

Note that (4), by definition of the Laplace transform [1], is an integral of the form

$$E(s) = \int_0^{\infty} \varepsilon(t) e^{-st} dt. \quad (5)$$

If e^{-st} is represented as an alternating series, then (5) can be written as

$$E(s) = \int_0^{\infty} \varepsilon(t) dt - s \int_0^{\infty} \frac{t \varepsilon(t)}{1!} dt + s^2 \int_0^{\infty} \frac{t^2 \varepsilon(t)}{2!} dt - \dots + s^i \int_0^{\infty} \frac{(-t)^i \varepsilon(t)}{i!} dt + \dots \quad (6)$$

Series (6) can also be represented as an infinite power series [2]

$$E(s) = \mu_0 + \mu_1 s + \mu_2 s^2 + \dots = \sum_{i=0}^{\infty} \mu_i s^i, \quad (7)$$

in which the weighting coefficients μ_i are written in the form

$$\mu_i = \int_0^{\infty} \frac{(-t)^i \varepsilon(t)}{i!} dt. \quad (8)$$

Assuming a stable object, for which the real parts of the roots of the characteristic equation of the mathematical model (1) are negative and are in the left half-plane of the complex plane, we can conclude that series (7) is convergent, and, consequently, the coefficients μ_i are finite.

To establish the relationship between the coefficients μ_i and S_i , we use equation (4), in which we perform substitutions for $E(s)$ from (7) and the inverse model $W_m^{-1}(s)$ from (2). First, we transform equation (4) to the form

$$(1 - sE(s))W_m^{-1}(s) = 1. \quad (9)$$

After performing the specified substitutions, we obtain the following expression

$$1 - \mu_0 s - \mu_1 s^2 - \mu_2 s^3 - \dots + S_1 s - \mu_0 S_1 s^2 - \mu_1 S_1 s^3 - \dots + S_2 s^2 - \mu_0 S_2 s^3 - \dots + S_3 s^3 = 1. \quad (10)$$

After reducing similar terms and equating the coefficients at the same powers (10) of the variable s , we can obtain S_k from (2) in the form of a recurrence relation in the form

$$S_k = \mu_{k-1} + \sum_{i=0}^{k-2} \mu_i S_{k-i-1}. \quad (11)$$

An algorithm for identifying the parameters of a mathematical model of type (1) is proposed:

1. Based on measurements of the transition function $h(t)$ of the mathematical model under study in the form of a set of discrete points on the interval $t \in [0, T_{max}]$, the values of $e(t)$ are determined in the form (4).
2. The mathematical model of the controlled object is sought in the form of a transfer function of the type (1), $b_i = 0$.
3. We find the coefficients μ_i by numerical integration of expression (8), $i = 0, 1, 2, \dots$, and we calculate the coefficients S_k by the recurrent expression (11), $k = 1, 2, 3, \dots$, the coefficients a_k according to expression (3), $k = 1, 2, 3, \dots$.
4. The transition function of the obtained model is investigated, and its adequacy is determined.

Conclusion

The proposed algorithm can be suitable for determining the initial values for the adjustable coefficients of the PID controller, one of the approaches to tuning that is proposed in (3).

1. Korn G. A. Mathematical handbook for scientists and engineers: definitions, theorems, and formulas for reference and review / G. A. Korn, and T. M. Korn, 2nd, enl. and rev. ed. New York: McGraw-Hill, 2000, 1130 p.
2. Abramowitz M., Handbook of mathematical functions with formulas, graphs, and mathematical tables / M. Abramowitz, I.A. Stegun, 10th printing, National Bureau of Standards Applied Mathematics Series, USA, Washington, 1972, 470 p.
3. Kucherov D. Simulation of PID-Regulator Tuning Using Genetic Algorithm on Multi-Criteria Objective Function For Controlling Unstable Objects / D. Kucherov, O. Poshyvailo, I. Myroshnychenko, Proceedings of ITS-2023: Information Technologies and Security, November 30, 2023, Kyiv, Ukraine (in publication)

COMPUTER MODEL OF TUNED TRANSISTOR OSCILLATOR TO STUDY ITS EXCITATION PROCESSES

Oleksandr Pliushch, ORCID ID: 0000-0001-5310-0660, Taras
Shevchenko National University of Kyiv, Kyiv, Ukraine,
oleksandr.pliushch@knu.ua

Bohdan Zhurakovskiy, ORCID ID: 0000-0003-3990-5205, National
Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic
Institute", Kyiv, Ukraine, zhurakovskiybyu@tk.kpi.ua

Maryan Kyryk, ORCID ID: 0000-0001-9156-9347, Lviv Polytechnic
National University, Lviv, Ukraine, marian.i.kyryk@lpnu.ua

Tuned transistor oscillators can be excited by either step-voltage that is formed when the power supply is turned on or by internal thermal noises that occupy wide spectrum. While the nature of the former is understandable and well-studied, how exactly the latter excite tuned transistor oscillator is not properly researched. In the paper, computer simulation methods are proposed to study these processes. Computer simulation model of a tuned transistor oscillator is designed with account of the thermal noises present in the generator. The computer model development procedure is performed in three steps. Firstly, a system of differential equations for the analyzed circuit is developed. Secondly, these equations are transformed into the system of equations of difference. Finally, the equations of difference are translated into code in Matlab programming environment. Noises with normal distribution are simulated in the base of the transistor in the model to account for its thermal noises. In the experiment, the model of the tuned transistor generator is excited by the noise with three different levels of power. It is shown that all three noise levels excite the tuned oscillator although the smaller the power, the later oscillator is excited and reaches its stationary state. The studied processes are presented in great detail using Matlab graphics tools, which help to understand the nature of the processes. Obtained results indicate validity of the tuned transistor oscillator model and its effectiveness for studying excitation modes. Designed model can be easily modified to carry out research in many other areas of the tuned oscillators field of study.

Keywords: Tuned transistor oscillators, oscillator excitation, thermal noise, differential equations, equations of difference, computer simulation, computer model

Introduction

Until recently, computer simulation techniques have not been widely used for studying operation and performance of tuned transistor oscillators. Among the methods and approaches usually deployed for this purpose and described in the literature are different analytical methods that can be, despite the seeming simplicity of the oscillators, very complex and require solutions of the system of non-linear differential equations. To resolve this problem, this paper proposes an approach to creating a computer simulation model in Matlab environment to study the underlying processes in those tuned transistor oscillators and evaluate performance parameters for different scenarios of their operation. In this paper the authors tend to apply the designed computer simulation model to detailed research as to how thermal noise with normal distribution can excite a tuned transistor oscillator.

Design of the computer model

Computer simulation model for a tuned transistor oscillator can be constructed through the number of steps. Firstly, the process starts with developing the system of differential equations that describes the signals in the different circuits of the tuned transistor oscillator. Then, this system of differential equations is transformed into the system of the equations of difference. Finally, the system of equations of difference is developed into Matlab code. Once the computer model is created, it can be tested and validated. Computer simulation models usually have many parameters that can be varied in a wide range. And this is where the main advantage of computer simulation lies; it creates many ways for researching tuned computer oscillators in different modes of operation and for versatile values of the circuit parameters.

Computer simulation results

To understand tuned oscillator excitation, computer simulation of the designed model was carried out in the Matlab environment for the noise power values of $1.0\text{e-}4$, $1.0\text{e-}6$, and $1.0\text{e-}8$ Watts.

Fig.1 illustrate capacitor voltage and inductance current in the tuned circuit of the oscillator for the noise power level $1.0\text{e-}4$.

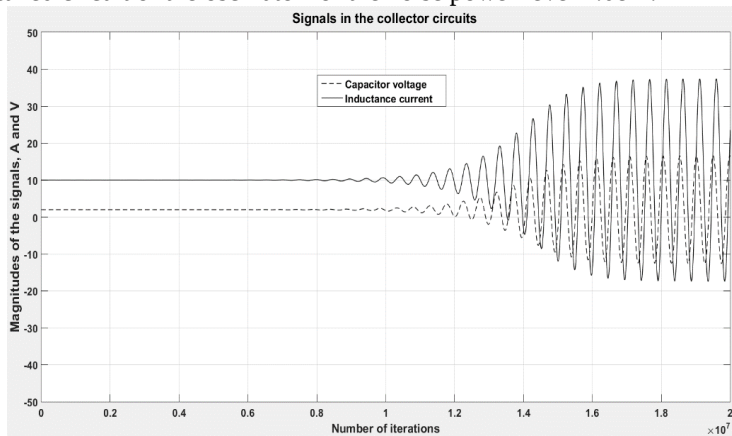


Figure1

As is seen from Fig.1, the thermal noise with normal distribution does effectively, as the theory states, start oscillations in tuned transistor oscillators. It is also obvious that certain time is required for the accumulation of the energy in the tuned loop before the threshold is reached and oscillations begin. In reality, this time of starting oscillations is not important with account of the overall purpose of the tuned oscillator.

Overall inference from the computer experiments with the designed computer simulation model allows one to stress that the model correctly represents operation of the tuned transistor oscillator and due to its simplicity and clarity, can be successfully used to study processes in the oscillators, as well as many others with some minor modifications.

Conclusion

This paper is devoted to studying whether or how thermal noise on its one can excite tuned transistor oscillator. To achieve this goal, an alternative approach to studying processes in oscillators has been proposed. This approach is based on the computer simulation model implemented in the Matlab environment. The developed model has been tested for the scenarios of the changing levels of the noise power in the oscillator. Simulation results demonstrated that the model correctly represents the operation of the oscillator and these results

comply with the general understanding of the processes in such devices. The main achievement is that it is proven that the thermal noises do excite tuned oscillators although differently for different levels of the thermal noise power. These results also support the conclusions and assumptions made in scientific literature by other authors. The computer simulation approach is simple, effective and can be modified to be used for a wide range of tuned oscillators. Designed computer model can be utilized for doing new promising studies in the field as well as improving performance of the existing devices.

1. O. Pliushch, Yu. Kravchenko, B. Zhurakovskiy, A. Savchenko. "Transmission Line Computer Simulation Model with Trapezoidal Rule." Proceedings of the IEEE International Scientific-Practical Conference PIC S&T, Kyiv, Ukraine, October 10–12, 2022, Paper 83, pp. 341–344. <https://doi.org/10.1109/PICST57299.2022.10238616GmbH>, 2021

2. O. Pliushch, Yu. Kravchenko, O. Matviichuk-Yudina, A. Rybydajlo, O. Trush, R. Mykolaichuk. "Computer model of radio frequency power amplifier." Proceedings of the IEEE 3rd International Conference on Advanced Trends in Information Theory, ATIT 2021, Kyiv, Ukraine, December 20–21, 2021, Paper 95, pp. 159–164. <https://doi.org/10.1109/ATIT54053.2021.9678673>

3. L. E. Frenzel Jr., Principles of electronic communication systems, 4th ed., McGraw-Hill Education, New York, NY, 2016.

4. I. M. Gottlieb, Practical Oscillator Handbook, Newnes, Jordan Hill, Oxford. 1997.

СИСТЕМА КЕРУВАННЯ ПРОЦЕСОМ ПАРОВО-КИСНЕВОЇ КОНВЕРСІЇ МЕТАНУ

Л. Ладісва, Б. Корнієнко, О. Пилипон

*Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»*

В даній роботі досліджувався процес конверсії метану. У виробництві водню і його сумішей найважливішим апаратом є конвертор, який і призначений безпосередньо для конвертування метану за участю кисню та парів води. Тому саме цей реактор досліджується, як технологічний об'єкт керування. Тиск в апараті підтримується витратою газів реакції на виході з конвертора і динамічні

характеристики цього параметру в подальшому не розглядаються. Також, не розглядаються втрати в навколишнє середовище, так як корпус конвертора був передбачений для цього і є теплоізоляція. Постійними є температури і концентрації вхідних потоків. Виходячи з наведеного вище, визначальним параметром даного процесу є концентрація метану на виході з конвертора. Щоб досягти заданої концентрації потрібно регулювати витратою кисню. При цьому забезпечується співвідношення потоків парогазової суміші і кисню. Витрати кисню входять в граничні умови математичної моделі конвертора як об'єкта з розподіленими параметрами. В даній роботі розроблена математична модель динаміки концентрації метану на виході з конвертора. Визначаються статичні та динамічні характеристики за каналами керування і збурення на основі створеної математичної моделі конвертора. Досліджується вплив допущень на вид і характер динамічних властивостей. Досліджувалася система в просторі стану. Запропонований критерій оптимальності. Знайдене оптимальне керування процесом конверсії метану. Синтезовано оптимальний лінійний регулятор. Даний підхід дозволив синтезувати оптимальний лінійний закон на основі застосування нелінійного диференційного рівняння Рікатті з розподіленими параметрами. Знайдено оптимальне керування процесом конверсії метану і оптимальна траєкторія переходу стану. Наведені графічні результати дослідження.

Ключові слова: математична модель; система; оптимальне керування; процес конверсії метану.

Актуальність пошуку оптимального керування процесу парокисневої конверсії метану полягає в тому, що до теперішнього часу не було запропоновано єдиного підходу до моделювання багатокомпонентних промислових процесів конверсії метану та не розв'язувалася задача на основі лінійного регулятора.

Каталітична конверсія природного газу нині стала основним методом отримання водню і синтезу газу для провідних галузей народного господарства. Найбільшими споживачами водню і його сумішей з окисом вуглецю або азотом є нафтопереробна (гідрогенізаційні процеси), хімічна і нафтохімічна (процеси

гідрування, синтез метанолу, бутанолу, вищих спиртів, аміаку, карбаміду, штучних палив, пластмас, синтетичних волокон тощо), харчова, енергетична, металургійна та інші галузі промисловості. Промисловими методами отримання водню та його сумішей конверсією природного газу є процеси парової, пароповітряної та парокисневої з подальшою конверсією окису вуглецю.

Паро-киснева конверсія метану є одним із важливих процесів в хімічній промисловості, оскільки вона дозволяє отримувати корисні хімічні продукти з природного газу. Процес полягає в реакції метану з паром води та киснем при певних температурах та тисках, що призводить до утворення синтез-газу та інших продуктів. Метан може бути використаний як джерело енергії або хімічної сировини для виробництва різноманітних сполук.

Аналіз останніх досліджень та публікацій в області реакційної інженерії та каталізу показує значний прогрес у розробці кінетичних моделей для опису реакційних механізмів та кінетики паро-кисневої конверсії метану. В даний час існує три методи окислювального перетворення метану: парціальне окислення, риформінг вуглекислотний та паровий риформінг метану та природного газу які споживають багато енергії в процесі підтримання реакції [1-3]. Розроблені математичні моделі кінетики риформінгу метану для каталітичних систем у реакторах різної конструкції [4-5]. Однак застосування цих моделей для керування процесом виявляється досить складним завдяки великій кількості параметрів та неоднорідності реальних умов.

З метою створення оптимальної системи керування процесом конверсії метану виникає необхідність у розробці математичної моделі, яка б дозволяла ефективно керувати процесом. Така модель повинна враховувати не лише кінетичні параметри реакції, але й динаміку системи, реакційні умови та параметри робочого середовища. На основі цієї моделі можна буде розробити алгоритми керування, які забезпечать оптимальну продуктивність та ефективність процесу конверсії метану.

Мета роботи - підвищення ефективності процесу паро-кисневої конверсії метану з використанням методу математичного моделювання для прогнозування роботи установки і оптимізації технологічного режиму в умовах мінливого складу сировини. Розробка і дослідження системи оптимального керування технологічного процесу.

Сучасні технології конверсії метану представлені широким класом різноманітних промислових процесів, які, незважаючи на всі відмінності, мають принципову спільність механізму.

Теперішні виробництва важко уявити без автоматизованої системи керування, що значно полегшує процес виробництва та зменшує кількість браку. Тому основною задачею даної роботи стало дослідження і створення системи керування процесом конверсії метану у рідкій фазі. Розроблена система керування повинна підтримувати задану концентрацію метану на виході з реактора що в свою чергу, повинно забезпечувати потрібну якість вихідного матеріалу та нормальне протікання процесу. В якості керуючого впливу було обрано витрату кисню, яка подається у певному співвідношенні відносно витрат парогазової суміші.

Аналіз попередніх робіт показав, що процеси конверсії метану та їх моделі в достатній мірі дослідженні. Та попри це, існуючі моделі не враховують особливості розподілу концентрації метану по висоті реактора, що не дозволяє дослідити систему керування. Тому основною задачею роботи стало розробити математичну модель процесу з розподіленими параметрами для дослідження системи керування з метою підвищення якості та енергозбереження. Керуючий вплив входить в граничні умови моделі.

На основі даної моделі процесу досліджувалася оптимальна система керування. Введено інтегральний квадратичний критерій якості. Даний підхід дозволив синтезувати оптимальний лінійний закон на основі розв'язання диференційного нелінійного рівняння Ріккати з розподіленими параметрами. Та розроблено алгоритм розрахунку оптимального керування процесом конверсії метану.

[1] Li Z., Chuang Y., Chen C.-T. et al. Modeling and optimization of methane steam reforming process for hydrogen production. International Journal of Hydrogen Energy. Vol. 46, No. 46. 25472-25483. (2021).

[2] Foo X. Y. Modeling and Control of Methane Conversion Processes. Elsevier, 300. (2023).

[3] Bartos R., Śmiechowski M., Libicki J. Modeling and simulation of methane conversion processes in fixed bed reactors. Chemical and Biochemical Engineering Quarterly, Vol. 37, No. 3, 261-271. (2021).

[4] Fini T., Patz C., Wentzel R. Oxidative Coupling of Methane to Ethylene: Senior Design Reports, 268. (2014).

[5] Korniyenko, B., Ladieva, L. Method of static optimization of the process of granulation of mineral fertilizers in the fluidized bed. Lecture Notes

ОЦІНКА ЗНАНЬ У ЗАКЛАДАХ ВИЩОЇ ОСВІТИ НА ОСНОВІ АДАПТИВНОЇ IRT-МОДЕЛІ

Сергій Любарський² [0000-0001-8068-1106, sergii.liubarskyi@viti.edu.ua],

Едуард Бовда² [0000-0002-8267-2120, edepig8305@ukr.net],

Юрій Самохвалов^{1,2} [0000-0001-5123-1288, yu1953@ukr.net]

¹*Київський національний університет імені Тараса Шевченка*
²*Військовий інститут телекомунікацій та інформатизації імені Героїв Крут*

Запропоновано підхід до оцінки рівня засвоєння навчального матеріалу здобувачами освіти. Цей підхід включає використання адаптивних IRT-моделей оцінки знань на основі однопараметричної моделі Раша, організацію *flexilevel*-процедури тестування та перевірку збалансованості тестових завдань. Це дає змогу визначити об'єктивний рівень підготовленості здобувача освіти за процедурою тестування та оцінити якість підготовки тестових завдань.

Ключові слова: Комп'ютерна система навчання, оцінювання рівня знань, адаптивне тестування, параметрична модель Раша, складність тестових завдань.

Актуальність

В теперішній час сучасні заклади вищої освіти розвинутих країн Європи активно працюють над вдосконаленням системи освіти, яка здатна інтегруватися у європейський освітній та науковий простори. Одним із рушійних механізмів такого вдосконалення є забезпечення справедливої конкуренції на ринку освітніх послуг із належним рівнем якості вищої освіти.

Процес підготовки високоякісних спеціалістів має передбачати застосування не тільки сучасних методів та засобів отримання нових знань, а і об'єктивну оцінку рівня їх засвоєння. Ефективність оцінювання залежить від дотриманням дидактичних принципів системності та об'єктивності.

Серед недоліків, що властиві традиційним підходам в оцінюванні знань здобувача освіти, слід відзначити низький рівень об'єктивності оцінки. Для забезпечення максимальної

інформативності та об'єктивності результатів контролю необхідно, щоб процедура тестування використовувала множини однорідно підібраних за рівнем складності тестових завдань.

За *flexilevel*-процедурою *адаптивного тестування* здобувачу освіти надається тестове завдання з множини середнього рівня складності і, залежно від відповіді на нього, наступне буде запропоновано легшим або складнішим за попереднє. Так відбувається до тих пір, поки не завершиться процедура тестування.

Метою дослідження є підвищення якості організації навчального процесу для здобувачів освіти через індивідуалізацію навчання та застосування гнучких методів визначення об'єктивного рівня засвоєння навчального матеріалу.

Основні положення

Під оцінюванням знань розуміється процедура дидактичного контролю, що спрямована на оцінку рівня освітньої підготовки здобувачів освіти, кількісну та якісну інтерпретацію показників досягнутих знань за допомогою відповідних шкал, критеріїв і норм. Цей підхід базується на інтеграції адаптивної моделі теорії Item Response Theory (ITR), заснованої на однопараметричній моделі Раша, та забезпеченням належного рівня збалансованості тесту.

Основним параметром завдань адаптивного тесту є рівень їх складності, що визначається експериментальним шляхом. Для вирішення даного завдання пропонується застосувати однопараметричну модель Раша. Вона використовує два основних параметри: θ – латентний параметр рівня знань учасника тестування і b – складність завдання.

У цій моделі використовується імовірність правильної відповіді на питання, яка обчислюється за формулою:

$$P_j(\theta) = \frac{e^{(\theta - b_j)}}{1 + e^{(\theta - b_j)}}.$$

Імовірність правильної відповіді на питання залежить тільки від різниці між рівнем знань учасника тестування і складністю завдання. За моделлю Раша тестові завдання не повинні розрізнятися за дискримінацією, вони можуть розташуватися тільки за рівнем їх складності.

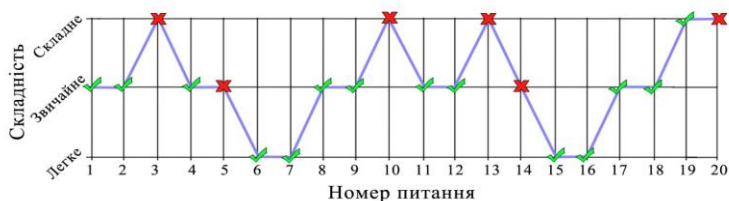
Оцінка знань на основі адаптивної *IRT*-моделі включає наступні кроки: підрахунок поточної успішності здобувача освіти

(процентне співвідношення балів, що надані на даному етапі проходження тесту до максимальної кількості балів за правильні відповіді); визначення поточної частки вірних і хибних відповідей; обчислення натурального логарифму відношення частки правильних відповідей на всі завдання тесту, до частки неправильних відповідей; обчислення складності завдання (частка неправильних відповідей на завдання тесту до частки всіх правильних відповідей на те ж завдання по множині здобувачів освіти); обчислення імовірності правильної відповіді здобувача освіти $P_f(\theta)$ на тестове завдання i , на основі її значення, визначення переходу на відповідний рівень складності завдання з встановленого експериментальним шляхом діапазону.

Порогове значення імовірності правильної відповіді здобувача освіти $P_f(\theta)=0,5$ виступає умовою зміни рівня складності завдання в процесі тестування. Рівень складності завдання збільшується при $P_f(\theta)\geq 0,5$ і зменшується у протилежному випадку.

Оптимальне значення показника диференціальної здатності тестових завдань забезпечується експериментально підібраними коефіцієнтами, збалансованість яких досягається за методом М. Челишкової [2].

Процес проходження тестування з поетапною зміною рівня складності можна представити наступним рисунком.



Визначення остаточної оцінки за традиційною системою оцінювання (незадовільно, достатньо, задовільно, добре, дуже добре, відмінно) здійснюється шляхом використання шкал вимірювань [2] виходячи із рівня успішності здобувача освіти.

Висновки

Запропоновано підхід до оцінки рівня засвоєння навчального матеріалу здобувачами освіти. Цей підхід включає використання адаптивних *IRT*-моделей оцінки знань на основі однопараметричної моделі Раша, організацію *flexilevel*-процедури

тестування та перевірку збалансованості тестових завдань. Оцінювання знань здійснюється генерацією питань з множин завдань різних рівнів складності для кореляції стратегії оцінювання. Оптимальна диференціююча здатність тестових завдань забезпечена експериментально підібраними коефіцієнтами, збалансованість яких досягається за методом М. Челишкової. Це дає змогу визначити об'єктивний рівень підготовленості здобувача освіти за процедурою тестування та оцінити якість підготовки тестових завдань.

Описаний підхід, порівняно з існуючими аналогами, надає змогу побудувати об'єктивний ряд переважності здобувачів освіти один над одним за результатами тестової оцінки, має високу наочність і помірну трудомісткість обчислень.

[1] Федорук П. І. Адаптивні тести: загальні положення. Математичні машини і системи, (1), 115–127 (2008). <http://dspace.nbuv.gov.ua/xmlui/bitstream/handle/123456789/748/10-fedoruk.pdf?sequence=1>.

[2] Анненкова І. П. Основи педагогічних вимірювань: навч.-метод. посіб. / І. П. Анненкова, Н. В. Кузнєцова, Л. А. Раскола – Одеса: Одес. нац. ун-т ім. І. І. Мечникова, 2021. – 210 с. ISBN 978-617-689-422-3

ГЕОМЕТРИЧНЕ МОДЕЛЮВАННЯ ФАНТАСТИЧНОЇ ЗБРОЇ ДЛЯ ОНЛАЙН КОМП'ЮТЕРНОЇ ГРИ У ЖАНРІ ШУТЕРУ

І. Куцербов, П. Ломовцев

*Одеський національний технологічний університет
(lomovtsevp@gmail.com)*

Гравці та розробники прагнуть створювати та використовувати оригінальні 3D-моделі, які додають іграм унікального стилю та підкреслюють їхню ідентичність. Такі моделі стають ключовою частиною гри, допомагаючи їй вирізнитися серед конкурентів. Особливим попитом користуються моделі зброї, мобів та інших об'єктів, які формують унікальний характер ігрового світу.

Сучасні програми, як-от Blender, ZBrush та Substance Painter, значно розширили можливості створення високоякісних моделей. Інструменти, які раніше були доступні лише великим студіям, тепер активно використовуються інді-розробниками. Це дає

змогу невеликим командам мінімізувати витрати, прискорити розробку та створювати конкурентоспроможні продукти, зокрема у жанрі Sci-Fi.

Інді-ігри все частіше інтегрують готові або кастомізовані моделі, що дозволяє ефективніше розподіляти ресурси. Завдяки платформам, таким як Unity Asset Store та Unreal Marketplace, розробники можуть купувати й продавати моделі, оптимізуючи витрати на розробку. Це стимулює розвиток ринку 3D-контенту, включаючи унікальні Sci-Fi моделі.

Сучасні технології ретопології й текстурювання забезпечують створення якісних моделей із оптимальним числом полігонів, що особливо важливо для онлайн-ігор. Процедурне текстурювання та рендеринг у реальному часі, наприклад за допомогою Marmoset Toolbag, дозволяють швидко оцінювати результати та знижувати витрати на доопрацювання.

Висока деталізація та інтеграція у рушії, як-от Unreal Engine і Unity, перетворюють унікальні моделі на ключові елементи для створення реалістичного ігрового досвіду. Моделі зброї та мобів, створені за стандартами PBR, забезпечують реалістичне відображення текстур і матеріалів, відповідно до сучасних вимог.

Процес роботи над 3D-моделями включає кілька етапів. Спочатку готуються всі необхідні ресурси: моделі, текстури, шейдери. Потім здійснюється тестування апаратних параметрів для аналізу продуктивності моделей у різних умовах. Оцінка якості охоплює реалістичність, деталізацію та відповідність естетичним критеріям. Ця оцінка може бути як суб'єктивною, так і об'єктивною, з використанням вимірювань.

Технічний аналіз продуктивності включає тестування моделей на різних конфігураціях комп'ютерних систем, оцінку швидкості рендерингу, споживання ресурсів і взаємодії з рушіями Unreal Engine чи Unity. Оптимізація спрямована на зниження системного навантаження без втрати якості. Додатково проводиться статистичний аналіз, який допомагає визначити найкращі варіанти текстур і моделей, будувати графіки продуктивності та порівнювати різні методи.

Заключним етапом є підготовка висновків, де узагальнюються результати роботи, формуються рекомендації щодо найефективніших методів моделювання та текстурювання, а також пропонуються оптимальні налаштування для різних систем. У висновках також враховуються можливості для

подальших досліджень та вдосконалення процесів створення 3D-асетів для ігор.

1. Blender Base Camp - Посібник з дизайну футуристичної зброї [Електронний ресурс] <https://www.blenderbasecamp.com/futuristic-weapons-blender-design-guide/>
2. Моделювання науково-фантастичної зброї для ігор у 3ds Max [Електронний ресурс] <https://www.pluralsight.com/courses/modeling-sci-fi-weapons-games-3ds-max-910>
3. Моделювання футуристичної зброї у Blender [Електронний ресурс] <https://www.blenderbasecamp.com/futuristic-weapons-blender-design-guide/>

АНАЛІЗ ТА ПОРІВНЯЛЬНЕ ДОСЛІДЖЕННЯ АЛГОРИТМІВ ТЕКСТУРУВАННЯ ДЛЯ РЕАЛІСТИЧНИХ ІГРОВИХ МОДЕЛЕЙ

Є. Загорій, П. Ломовцев

*Одеський національний технологічний університет
lomovtsevp@gmail.com*

Основні методи текстурювання реалістичних моделей діляться на ручне текстурювання (Hand-paint), PBR (Physically based render), процедурне текстурювання. Кожен із цих методів має свої сфери застосування, але в межах розробки ігор ці сфери відчутно звужуються до певних жанрів ігор, а також залежать від пристроїв, які буде застосовано для гри.

Текстурювання вручну - це процес створення текстур вручну, часто з використанням цифрових інструментів для малювання, таких як Photoshop, Procreate або спеціалізованого 3D-програмування, наприклад, Substance Painter. Ця техніка підкреслює художню творчість, коли художники малюють деталі безпосередньо на моделях, надаючи їм унікального, стилізованого і часто дуже деталізованого вигляду.

Художники створюють текстури з нуля, промальовуючи такі деталі, як затінення, колірні варіації, вивітрування, знос і дефекти поверхні. Ручне малювання дозволяє детально кастомізувати та стилістично виражати, що корисно для досягнення особливого вигляду, особливо для стилізованих або художніх ігрових ресурсів. На відміну від процедурних або автоматизованих методів, ручне малювання фокусується на

складних деталях, які сприяють більш персоналізованому та візуально переконливому результату. Цей метод популярний для стилізованих ігор та 2D/3D ілюстрацій, де метою часто є посилення художньої виразності, а не досягнення гіперреалістичних візуальних ефектів.

Текстурування PBR (Physically-Based Rendering) - це техніка, спрямована на створення більш реалістичних і фізично точних матеріалів для 3D-моделей. Основний принцип PBR полягає в імітації взаємодії світла з поверхнями, в результаті чого матеріали виглядають більш реалістично і послідовно при різних умовах освітлення.

PBR-текстури визначають матеріали за допомогою таких властивостей, як шорсткість, металевість і базовий колір. Ці властивості ґрунтуються на реальній фізиці, забезпечуючи точні відображення, освітлення та затінення. Робочі процеси PBR включають створення базового кольору, шорсткості, металевих, нормальних і навколишніх карт оклюзії. Ці карти описують, як поверхні реагують на світло, і впливають на загальний вигляд матеріалу. На відміну від традиційних методів текстурування, PBR забезпечує більш реалістичну взаємодію затінення і освітлення, що допомагає досягти більш захоплюючого візуального досвіду, особливо в іграх і CGI. PBR став стандартом в індустрії, особливо в розробці ігор та кіновиробництві, оскільки він забезпечує узгодженість на різних платформах та освітлювальних установках.

Процедурне текстурування - це створення текстур за допомогою алгоритмів замість того, щоб малювати кожен текстур вручну. Цей підхід спирається на алгоритми, які генерують візерунки та матеріали на основі таких параметрів, як шум, фрактали, математичні функції та інші процедурні методи. На відміну від традиційних текстур, намальованих вручну, процедурне текстурування не покладається на заздалегідь визначені зображення або мазки пензля, нанесені вручну, що дозволяє отримати більш динамічні та різноманітні результати.

Процедурні текстури можуть змінюватися в реальному часі, залежно від використовуваних параметрів. Вони гнучкі та адаптуються до різних типів поверхонь, умов освітлення та масштабу. Процедурна генерація текстур дозволяє дизайнерам економити час на створенні складних моделей і матеріалів, особливо для великих або повторюваних об'єктів в іграх або візуальних ефектах. Процедурні текстури можна легко

використовувати повторно і змінювати, оскільки вони генеруються на основі математичних або алгоритмічних правил, що зменшує потребу у створенні унікальних текстур для кожного ресурсу. Дизайнери можуть контролювати і налаштовувати процедурні параметри для створення різних варіацій текстури, забезпечуючи послідовний, але унікальний вигляд ресурсів. Процедурне текстурування широко використовується в іграх, де великомасштабні середовища з великою кількістю об'єктів потребують узгоджених і масштабованих текстур, наприклад, у іграх з відкритим світом або науково-фантастичних фільмах.

АВТОМАТИЗОВАНИЙ ПІДХІД ДО ФОРМУВАННЯ ЖУРІ СТУДЕНТСЬКИХ НАУКОВИХ КОНКУРСІВ НА ОСНОВІ НАУКОМЕТРИЧНИХ ТА АЛЬТМЕТРИЧНИХ ПОКАЗНИКІВ

**Віталій Циганок^{1,2}, Ярослав Хроленко¹, Ірина Доманецька²,
Ольга Циганок³, Олена Трубіцина⁴**

*¹Інститут проблем реєстрації інформації НАН України²
Київський національний університет імені Тараса Шевченка*

³Державний торговельно-економічний університет

⁴Київський національний лінгвістичний університет

tsyganok@ipri.kiev.ua; yarrik40.000@gmail.com;

domanetska@knu.ua; olzyg@ukr.net; elenatrubitsyna88@gmail.com

Розглянуто автоматизований підхід до формування складу журі для оцінювання наукових робіт студентів. Запропоновано використання наукометричних та альтметричних показників як критеріїв оцінювання кандидатів. Наукометричні показники, такі як h-індекс, кількість публікацій та цитувань, дозволяють оцінити академічну продуктивність дослідників, тоді як альтметрики враховують цифровий вплив і суспільну активність. Для вирішення задачі відбору запропоновано багатокритеріальну модель, яка використовує метод TOPSIS. Автоматизація збору та аналізу даних базується на інтеграції з платформами Scopus, Web of Science, Altmetric.com тощо. Розроблений підхід сприяє об'єктивності, прозорості та ефективності оцінювання, що підвищує якість проведення конкурсів. Перспективи

розвитку включають вдосконалення моделей вагових коефіцієнтів та подальшу автоматизацію процесів.

Ключові слова: конкурси наукових робіт, формування журі, наукометричні показники, альтметрики, багатокритеріальна модель, метод TOPSIS

Вступ

Конкурси наукових робіт студентів є важливим елементом системи вищої освіти, що сприяє розвитку молодих науковців, їх інтеграції в академічну спільноту, а також стимулює дослідницьку активність і творчість[1]. Успіх таких заходів значною мірою залежить від складу журі, яке оцінює роботи учасників. Формування журі має базуватися на чітких критеріях, що гарантують об'єктивність, неупередженість і високий професійний рівень оцінювання. У цьому дослідженні запропоновано підхід до формування складу журі, який базується на використанні наукометричних та альтметричних показників, а також автоматизації процесу формування журі.

Виклад основного матеріалу

Задля забезпечення збалансованого підходу до відбору членів журі пропонується у якості критеріїв оцінювання експертів обертинукометричні та альтметричні показників публікаційної активності експертів-кандидатів.

Класичні наукометричні показники, такі як h-індекс, кількість публікацій та цитувань, є перевіреними інструментами для оцінки академічної продуктивності. H-індекс, наприклад, дозволяє оцінити науковий вплив дослідника, враховуючи як обсяг, так і цитованість його робіт[2]. Високий h-індекс свідчить про стабільну наукову продуктивність і визнання в академічній спільноті. Однак цей показник має певні недоліки, такі як чутливість до самоцитування та обмеження у відображенні впливу нових робіт. Доповнення його кількістю публікацій та цитувань дозволяє отримати повнішу картину наукової активності кандидата.

Однак сучасна наукова комунікація виходить за межі традиційних академічних рамок, тому використання альтметричних показників стає важливим доповненням. Альтметрики враховують цифровий слід досліджень та їхній суспільний вплив. Наприклад, AltmetricScore відображає

активність у соціальних медіа, блогах та новинах, що дає змогу оцінити реакцію суспільства на наукову роботу [3]. Це особливо корисно для оцінки актуальності досліджень та їхньої популярності серед широкої аудиторії. Міждисциплінарність також є важливим критерієм, оскільки вказує на здатність експерта оцінювати роботи, що охоплюють різні галузі знань. Такі навички є особливо актуальними для міждисциплінарних конкурсів, де потрібно оцінювати роботи, що поєднують кілька наукових напрямів. Соціальна активність експерта, зокрема його присутність на платформах як-от ResearchGate чи LinkedIn, свідчить про його здатність популяризувати науку та комунікувати зі студентами та громадськістю. Це важливо для конкурсу, де члени журі мають не лише оцінювати роботи, а й надавати конструктивний зворотний зв'язок.

Додатково, організаційні критерії, такі як відсутність конфлікту інтересів, рецензентський досвід і професійне визнання, допомагають забезпечити об'єктивність оцінювання. Незалежність експертів та їхній досвід у рецензуванні гарантують неупередженість та дотримання академічних стандартів, що є критично важливим для чесного проведення конкурсу. Наявність нагород та участь у наукових комітетах чи редакційних радах підкреслюють високий рівень експертності.

Для вирішення задачі формування складу журі запропоновано використовувати багатокритеріальну модель прийняття рішень. Кандидати оцінюються за зваженими показниками академічного впливу, цифрового сліду та організаційних критеріїв. Мета - максимізувати загальну зважену суму оцінок обраних експертів

$$\text{Maximize } \sum_{i=1}^m (\alpha \cdot h_index_i^{\text{норм}} + \beta \cdot \text{Altmetric_Score}_i^{\text{норм}} + \gamma \cdot \text{Org_Links}_i^{\text{норм}}) \cdot y_i$$

де

$y_i \in \{0,1\}$ — це бінарна змінна, яка дорівнює 1, якщо кандидат i обраний до журі, і 0 в іншому випадку. m - потужність експертної групи (журі).

$h_index_i^{\text{норм}}$, $\text{Altmetric_Score}_i^{\text{норм}}$, $\text{Org_Links}_i^{\text{норм}}$ - унормовані значення наукометричних, альтметричних показників та показників, що обумовлюються умовами конкурсу, для того експерта.

Модель базується на методі TOPSIS (Technique for Order of Preference by Similarity to Ideal Solution), що

дозволяє обрати експертів, найближчих до ідеального рішення[4]. Для нормалізації та агрегування даних пропонується використовувати автоматизовані інструменти, інтегровані з базами даних Scopus, Web of Science, Altmetric.com, а також соціальними платформами ResearchGate і LinkedIn.

Автоматизація процесу збору та аналізу даних забезпечує швидкість, прозорість і мінімізацію впливу людського фактора. Розроблена логічна архітектура застосунку включає користувацький інтерфейс для зручного доступу до метрик, серверну частину для обробки даних і базу даних для збереження результатів. Зручний інтерфейс дозволяє порівнювати кандидатів, візуалізувати метрики і здійснювати пошук за заданими параметрами.

Висновки

Запропонований підхід сприяє покращенню об'єктивності та ефективності формування складу журі, що в свою чергу підвищує якість проведення конкурсів. Перспективи розвитку дослідження включають розробку математичних моделей для визначення оптимальних вагових коефіцієнтів у багаторівневій системі оцінювання. Автоматизація відбору експертів є важливим кроком до вдосконалення організації студентських наукових конкурсів, сприяючи формуванню нових стандартів академічної справедливості.

1. V. Tsyganok, Y. Khrolenko, and I. Domanetska, "Information Technology for Student Scientific Works Competitions," in *Proc. 1st Int. Student Conf.: Digital Generation*, Astana, Kazakhstan, Apr. 2024, pp. 310-315. Available: <https://student-conference.astanait.edu.kz/>
2. Hirsch J.E. An index to quantify an individual's scientific research output, *Proceedings of the National Academy of Sciences*, 2005. DOI: 10.1073/pnas.0507655102.
3. Altmetrics: a manifesto [Електронний ресурс]. Доступ: <https://altmetrics.org/manifesto>.
4. Thakkar, Jitesh. (2021). Technique for Order Preference and Similarity to Ideal Solution (TOPSIS). DOI: 10.1007/978-981-33-4745-8_5.

ПЛАНУВАННЯ МІСІЙ БПЛА НА ОСНОВІ ОНТОЛОГІЙ ТА РОЛЬОВОГО РОЗПОДІЛУ ЗАВДАНЬ У МУЛЬТИАГЕНТНІЙ СИСТЕМІ

Анатолій Гладун, ORCID: 0000-0002-4133-8169,
glanat@yahoo.com,

Катерина Хала, ORCID: 0000-0002-9477-970X,
cecerongreat@ukr.net,

Олександр Волков, ORCID: 0000-0002-5418-6723,
alexvolk@ukr.net

Інститут інформаційних технологій та систем НАН України

Розглядається актуальна проблема планування завдань у МАС рою БПЛА, яка спрямована на підвищення ефективності виконання місій в умовах обмежених ресурсів та динамічного середовища. Запропоновано онтологічний підхід до розподілення ролей і завдань, який враховує пріоритети місії, логічну ієрархію підзавдань та обмеження середовища, забезпечуючи координацію та автономність агентів у реальному часі.

Ключові слова: МАС, рій БПЛА, онтологія, планування завдань, розподілення ролей.

Вступ

Проблема планування завдань у мультиагентній системі (МАС) рою безпілотних літальних апаратів (БПЛА) є актуальним завданням, спрямованим на підвищення ефективності виконання місій в умовах обмежених ресурсів, динамічних середовищ та необхідності координації агентів для досягнення спільної мети. Водночас складність децентралізованого керування роєм БПЛА полягає у забезпеченні злагодженої взаємодії між агентами без централізованого контролю, з урахуванням обмежень обчислювальних ресурсів, затримок обмінювання даними та мінливості зовнішнього середовища [1].

Розподілення ролей та завдань відповідно до місії

Агент є інструментом керування БПЛА та перебирає на себе його можливості, забезпечуючи автономність у прийнятті рішень, обробці інформації та виконанні завдань відповідно до заданих ролей і пріоритетів місії. Наприклад, у місіях реагування на

катастрофи, таких як оцінювання та пом'якшення наслідків поведінки в міській місцевості [2].

Під місією будемо розуміти клас на вершині онтології, який описує загальне завдання призначене рою БпЛА. Місія охоплює низку завдань, які далі поділяються на підзавдання та мають розподілятися між агентами БпЛА відповідно до їхніх ролей. Розбиття завдань на підзавдання розглядається на основі методу комбінаторної оптимізації [3]. Ролі агентів є наперед визначеними в онтології та розподіляються між агентами залежно від типів і можливостей БпЛА, якими вони керують. Складові структури для місії реагування на катастрофу описано в таблиці 1. Агенти можуть мати кілька ролей під час однієї місії, але на початковому етапі за ними закріплюється та, яка має максимальну відповідність характеристикам БпЛА та найвищий пріоритет у місії.

Таблиця 1 – Приклад складових структури місії

Завдання	Підзавдання	Пріоритет	Необхідна роль
аерофотокартування	отримання зображень	високий	картограф
	зшивання зображень	середній	картограф
	3d моделювання	середній	картограф
пошук і порятунок	візуальне сканування	високий	розвідник
	теплове сканування	високий	розвідник
	позначення місця розташування	високий	рятувальник
доставка ресурсів	завантаження запасів	середній	транспортер
	доставка	середній	транспортер

Онтологічний підхід до розподілення ролей і завдань

Онтологія структурує ролі агентів БПЛА, завдання та підзавдання, формуючи ієрархію, що відображає зв'язки між завданнями, профілями агентів, пріоритетами місії та обмеженнями середовища.

Процес розподілення ролей охоплює 1) зіставлення онтологій, де онтологічний опис можливостей БпЛА (яким керує агент) зіставляється з ролями зафіксованими в онтології, та

підбирається роль з найвищим значенням ваг (ролей може бути кілька) [4]; 2) пріоритетне призначення, за якого домінуючими визначаються ролі з найвищими вагами та пріоритетами, а інші стають допоміжними. У таблиці 2 наведено розподілення ролей агентам відповідно до можливостей БпЛА.

Таблиця 2 – Розподілення ролей агентам відповідно до можливостей БпЛА згідно місії

БпЛА	Можливості	Призначена роль
UAV_1	зображення високої роздільної здатності, середня швидкість	картограф
UAV_2	зображення високої роздільної здатності, велика висота	картограф
UAV_3	високошвидкісна теплова камера	розвідник
UAV_4	висока швидкість, висока витривалість	розвідник
UAV_5	гучномовець, легкий, швидкий	рятувальник
UAV_6	медична аптечка, середньої витривалості	рятувальник
UAV_7	велика вантажопідйомність, середня швидкість	транспортер
UAV_8	висока вантажопідйомність, довговічність	транспортер
UAV_9	універсальний, середні можливості	картограф/розвідник
UAV_10	легка теплова камера малого радіусу дії	розвідник

Процес розподілення завдань охоплює 1) *зіставлення онтологій*, де онтологічний опис завдань чи підзавдань місії зіставляється з ролями, після чого агенти з підходящими ролями конкурують за завдання; 2) *аналіз відношень завдань/підзавдань*, який забезпечує дотримання логічного порядку їх виконання; 3) *динамічне розподілення підзавдань*, де у разі завершення завдання, зміни пріоритету або виходу БпЛА з ладу, агенти можуть конкурувати за інше підзавдань відповідно до їх другорядних ролей, онтологія оновлює відношення, і завдання перерозподіляються в режимі реального часу для забезпечити безперервність місії.

Оптимальне призначення завдань та підзавдань виконується за такою формулою:

$\max\{\sigma_1 \sum_{i=1}^n \sum_{j=1}^m p_j r_{ij} x_{ij} - \sigma_2 \sum_{i=1}^n \sum_{j=1}^m c_{ij} x_{ij}\}$, за умов:

$\sum_{j=1}^m x_{ij} \leq 1, \forall i = 1, 2, \dots, n$, БПЛА виконує не більше одного підзавдання та

$\sum_{i=1}^n x_{ij} \leq 1, \forall j = 1, 2, \dots, m$, кожне підзавдання виконується хоча б одним БПЛА,

де x_{ij} – змінна виконання j -го завдання i -м БПЛА, $x_{ij} \in \{0, 1\}$; p_j – пріоритет завдання j ; r_{ij} – частка завершення j -го завдання i -м БПЛА, $0 \leq r_{ij} \leq 1$; c_{ij} – затрати виконання j -го завдання i -м БПЛА; σ_1 і σ_2 – вагові коефіцієнти; n – кількість БПЛА; m – кількість завдань/підзавдань.

Висновок

Розглянуто онтологічний підхід до планування завдань у МАС рою БПЛА, який забезпечує ієрархічну структуру місії, врахування ролей агентів, пріоритетів завдань і зовнішніх обмежень. Запропонована методика дозволяє оптимально розподіляти ресурси та забезпечувати автономність агентів у виконанні завдань. Використання онтології для формалізації завдань і підзавдань, а також динамічний перерозподіл ролей у режимі реального часу забезпечують безперервність виконання місії навіть за умов змінних зовнішніх факторів чи виходу агентів БПЛА з ладу.

1. X. Wang, H. Wang, H. Zhang, M. Wang, A mini review on UAV mission planning, *J. of Industrial and Management Optimization*, Vol 19, iss. 5, 2023, pp. 1-21. DOI:10.3934/jimo.2022089.

2. О. Волков, М. Комар, Д. Волошенюк та ін., Інтелектуальний контроль, локалізація та картографування в геоінформаційних системах на основі аналізу візуальних даних, *Cybernetics and Computer Engineering*. 2020, № 200, с. 41–58.

3. Н. Тимофієва, Семантичне моделювання та комбінаторна оптимізація, *Вісник Херсонського національного технічного університету*, № 3(2). 2018, с. 101-106.

4. A. Gladun, K. Khala, R. Martinez-Bejar, Development of Object's Structured Information Field with Specific Properties for Its Semantic Model Building, in: *CEUR Workshop Proceedings*, Vol. 3241, 2021, pp. 102–111.

ВИКОРИСТАННЯ ТЕХНОЛОГІЙ ГІШ ДЛЯ АНАЛІЗУ КАСКАДНИХ ЕФЕКТІВ КРИТИЧНОЇ ІНФРАСТРУКТУРИ

А. Бойченко, ORCID: 0009-0003-2981-7209

*Інститут проблем реєстрації інформації НАН України
вул. М. Шпака, 2, 03113 Київ, Україна*

Постановка проблеми

В умовах безперервної російської агресії вчасний аналіз та передбачення каскадних ефектів (КЕ) викликаних атаками на об'єкти критичної інфраструктури (ОКІ) України вимагає використання найсучасніших методів та технологій.

Аналіз показує що взаємозалежність між об'єктами КІ може призвести до серйозних економічних і фізичних збитків у регіональному, загальнодержавному та навіть у планетарному масштабі. Виявлення та аналіз явних та прихованих зв'язків та залежностей, які існують між елементами КІ є необхідною складовою для розробки стратегій захисту та реагування на загрози живучості та резильєнтності КІ, а також мінімізації руйнівних наслідків деструктивних впливів [1].

Надзвичайна важливість забезпечення надійної роботи критичної інфраструктури для самого існування кожного окремого члена суспільства так і цивілізації у цілому спричинило справжню лавину публікацій у цій галузі.

У роботі [2] показано, що вразливість критичної інфраструктури може орієнтувати хід подій у зв'язку зі зворотним зв'язком суспільства, а не бути просто наслідком природних тригерів. Наші висновки вказують на зміну парадигми на етапі готовності, яка може включати точки ескалації та соціальні вузли, але це також розкриває абсолютно нову сферу досліджень для дослідників катастроф.

Багато досліджень КЕ доводять, що початкова подія складається з ланцюга можливостей, які в якийсь момент включають «головну подію» або найбільш катастрофічний наслідок. Основне припущення в цій моделі полягає в тому, що існує ланцюг слабких місць, у якому лінійна послідовність подій має односпрямований причинно-наслідковий зв'язок, і всі випадки чи інциденти є наслідками, а подія, що викликає, є основною причиною. Більшість критично важливих інфраструктур є надто складними, щоб ця модель могла нормально функціонувати [3 -

5]. Стаття [6] присвячена проблемі виявлення зв'язків для складних, непересічних мереж КІ, у яких існує високий ступінь взаємної взаємозалежності.

Дослідження [7] зосереджено на оцінці каскадної катастрофи і встановленні ступеня причинного зв'язку між елементами КІ. Виникає КЕ, наявність лінії причинно-наслідкового зв'язку між різноманітними інцидентами та їхнім впливом означає КЕ. Складність і рівень незалежності сучасного життя такі, що ймовірність КЕ у КІ стрімко зростає.

У публікаціях [8,9] показано, що ризик можна контролювати за допомогою підходу ризик-винагорода. Коли користувач вчиняє зловмисну дію, його рейтинг довіри знижується, поки доступ не буде повністю заборонено. Коли користувач бере участь у звичайному ви-користанні, його рейтинг довіри відновлюється.

У роботах [10, 11] запропоновано використання систем для генеративного штучного інтелекту (ГШІ) для формування семантичних моделей предметних областей.

Мета дослідження

Метою даного дослідження є розробка та дослідження методики використання технологій ГШІ для аналізу критичних інфраструктур. Ця методика включає засоби інтелектуального опрацювання текстових даних, побудови мереж взаємозалежності ключових понять, і врешті решт формування сценаріїв інформаційного впливу.

Результати дослідження

Головною складністю покращення якості мережевих моделей є відсутність достатньої вибірки даних (експертних оцінок) для призначення чисельних змінних, що описують процеси КІ і які саме й дозволяють налаштувати модель КЕ та зробити її більш реалістичною. Для подолання цього недоліку можливим є використання запропонованої у [10] методики формування рою віртуальних експертів, який дозволяє отримати дані не від живих фахівців а від доступних на даний час Інтернет систем:

- Character [12];
- Deepseek[13];
- Chatgpt_ca[14];
- ChatGPT Openai[15];
- Claude[16];

- Copilot [17];
- Gemini [18];
- GptFree [19];
- Groq[20];
- Zerogptai[21].

На рисунку 1 представлені приклади відповіді на один і той же запит двох систем ГШІ, Claude[16] та DeepSeek[13].

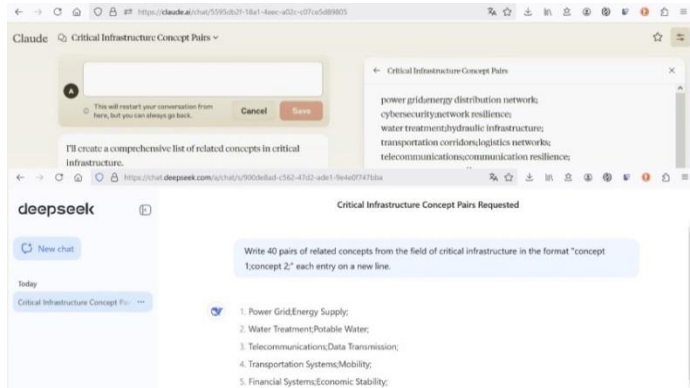


Рис. 1 Зразки відповідей від різних віртуальних експертів.

У роботі [1] нами було запропоновано підхід до моделювання сценаріїв за допомогою знання-орієнтованої технології з використанням графової БД Neo4j. Ця технологія включає зокрема механізми формування графових моделей з даних, що містяться в csv файлах. Отримані дані у форматі CSV заносяться до бази даних Neo4j і використовуються для моделювання.

Використання Neo4j дає можливість використовувати вже існуючі інструментальні засоби побудови моделей та метрики аналізу графових структур, які суттєво прискорюють процес дослідження розвитку KE. За допомогою цих засобів можна оцінювати поведінку KI при виникненні різних сценаріїв розвитку KE та виявити найбільш вразливі вузли KI, що впливає на прийняття рішень [22].

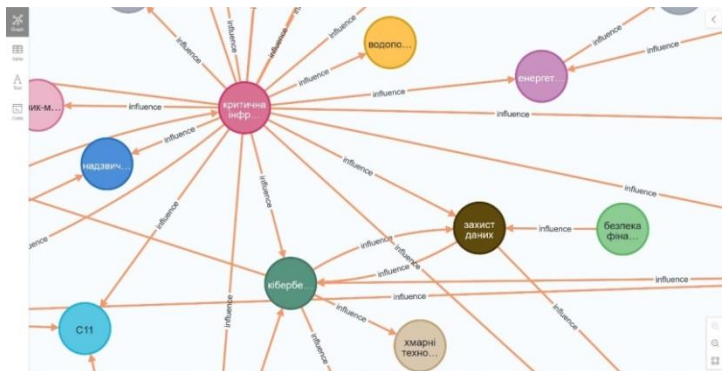


Рис. 2 Відображення у браузері фрагменту отриманої графічної бази даних.

Висновки

Запропонована методика використання технологій ГШІ для аналізу критичних інфраструктур, яка ключає засоби інтелектуального опрацювання текстових даних, побудови мереж взаємозалежності ключових понять, і врешті решт формування сценаріїв інформаційного впливу. Використання даного підходу дозволяє здійснювати аналіз каскадних ефектів пов'язаних критичних інфраструктур на базі моделюючого комплексу Інституту проблем реєстрації інформації НАН України. Результати моделювання, мають дозволити приймати більш обґрунтовані управлінські рішення.

[1]. Бойченко, А. В., & Сенченко, В. Р. (2022). Підхід до моделювання геопросторових каскадних ефектів критичних інфраструктур. Інформаційні технології та безпека, 56(7), 35.

[2]. Davydiuk, A., & Zubok, V. (2023, May). Analytical review of the resilience of Ukraine's critical energy infrastructure to cyber threats in times of War. In 2023 15th International Conference on Cyber Conflict: Meeting Reality (CyCon) (pp. 121-139). IEEE.

[3]. De Oliveira, A. K. B., Battemarco, B. P., Barbaro, G., Gomes, M. V. R., Cabral, F. M., de Oliveira Pereira Bezerra, R., ... & Miguez, M. G. (2022). Evaluating the role of urban drainage flaws in triggering cascading effects on critical infrastructure, affecting urban resilience. Infrastructures, 7(11), 153.

[4]. Lewis, L. P., & Petit, F. (2019). Critical infrastructure interdependency analysis: operationalising resilience strategies. Contributing

Paper to the Global Assessment Report on Disaster Risk Reduction (GAR 2019), 33pp.

[5]. Sun, W., Bocchini, P., & Davison, B. D. (2022). Overview of interdependency models of critical infrastructure for resilience assessment. *Natural Hazards Review*, 23(1), 04021058.

[6]. Aebi, S., Hauri, A., & Kamberaj, J. (2024). Critical Infrastructure Resilience in Ukraine: Energy, Transportation, and Communication. *CSS Risk and Resilience Reports*.

[7]. Menashri, H., & Baram, G. (2015). Critical infrastructures and their interdependence in a cyber attack—the case of the US. *Military and strategic Affairs*, 7(1), 22.

[8]. Barquet, K., Englund, M., Inga, K., André, K., & Segnestam, L. (2024). Conceptualising multiple hazards and cascading effects on critical infrastructures. *Disasters*, 48(1), e12591.

[9]. Abedi, A., Gaudard, L., & Romero, F. (2019). Review of major approaches to analyze vulnerability in power system. *Reliability Engineering & System Safety*, 183, 153-172.

[10]. Д.В. Ланде. OSINT у кібербезпеці : навч. пос. //OSINT у кібербезпеці : навч. пос. - Київ: ТОВ "Інжиніринг", 2024. - 522 с.

[11]. Lande, D., Strashnoy, L., Driamov, O., & Feher, A. Формування сценаріїв діяльності на базі сервісів генеративного штучного інтелекту. *Штучний інтелект*, 2023; 97(3): 94-103. ISSN 2710 - 1673. DOI: <https://doi.org/10.15407/jai2022.01.08>.

[12]. <https://character.ai>.

[13]. <https://chat.deepseek.com>.

[14]. <https://www.chatgpt.ca>.

[15]. <https://chat.openai.com>.

[16]. <https://claude.ai>.

[17]. <https://copilot.microsoft.com>.

[18]. <https://gemini.google.com>.

[19]. <https://gptfree.co>.

[20]. <https://groq.com>.

[21]. <https://zerogptai.org>.

[22]. Сенченко, В. Р., Бойченко, А. В., Коваль, О. В., Бисько, Р. М., & Хоменко, О. М. (2024). Огляд методів і технологій сценарного аналізу каскадних ефектів. *Реєстрація, зберігання і обробка даних*, 26(1), 24-54.

МОДЕЛЮВАННЯ НЕПОВНИХ ЕКСПЕРТНИХ ОЦІНОК, ЩО ВРАХОВУЮТЬ РАНЖИРУВАННЯ

Каденко С.В.¹, Андрійчук О.В.^{1,2}, Гоменюк Г.А.¹

¹ Інститут проблем реєстрації інформації НАН України,

² Київський національний університет імені Тараса Шевченка
*seriga2009@gmail.com; andriichuk@ipri.kiev.ua;
galina.gomeniuk@gmail.com*

Вступ

Доповідь ілюструє поточний етап дослідження, спрямованого на підвищення якості рішень, що приймаються у слабко структурованих областях, таких, як, наприклад, безпекове середовище. Рішення, що приймаються у таких областях почасти визначаються критеріями, факторами та альтернативами, для яких відсутні еталонні значення та одиниці виміру. Відтак, експертні оцінки, зокрема, представлені у вигляді парних порівнянь (ПП) та ранжирувань об'єктів та критеріїв, за якими вони оцінюються, є важливим джерелом інформації.

Для відносного зважування поб'єктів необхідно здійснити відп – 1 доп(п – 1)/2 їхніх ПП. Якщо кількість об'єктів та/або критеріїв їхньої оцінки є великою, то актуальною стає задача скорочення кількості необхідних ПП без втрати достовірності результатами експертизи. Одним зі шляхів розв'язання задачі є використання попередніх ранжирувань об'єктів як додаткового джерела інформації. Це дозволяє уникнути необхідності введення усієї множини ПП.

З-поміж методів ПП, що враховують ранжирування, можна виокремити метод best-worst[1], Top 2 (best-second best) [2] та метод черг[3], яким присвячено низку нещодавніх досліджень та публікацій (наприклад, [4]).

Поточні задачі

Поточні дослідження авторів спрямовані на порівняння трьох вказаних підходів за якістю з метою визначення найефективнішого з них. У якості критерію порівняння методів пропонується обрати їхню стійкість до помилок експертів (адже, внаслідок вищевказаних особливостей слабко-структурованих предметних областей, використання абсолютних показників, таких як точність, неможливе). Порівняння методів за стійкістю пропонується здійснювати у ході експерименту, в якому

багаторазово моделюватиметься весь цикл експертизи (без залучення реальних експертів). Авторами проведено численні експерименти подібного типу для оцінки якості інших методів підтримки прийняття рішень (наприклад [5]).

Моделювання експериментального дослідження з порівняння методів експертного оцінювання

Методи Top 2 та best-worst є неповними та, за означенням, вимагають від експерта здійснення рівно $(2n - 3)$ ПП поб'єктів. Метод черг може бути як повним, так і неповним [3]. Для забезпечення справедливого порівняння трьох вказаних методів необхідно подбати про те, щоб кількість ПП, які моделюються у них була однаковою. Для досягнення зв'язності множиною ПП у методі черг, їхня кількість має дорівнювати

$$N_{зв} = \begin{cases} \left(\binom{\frac{n-1}{2}}{4} + 1 \right), n \text{ не парне}, n > 1 \\ \left(\binom{\frac{n}{2}}{4} + 1 \right), n \text{ парне}, n > 2 \end{cases} \quad (1)$$

Відтак, для порівняння трьох методів має виконуватися нерівність $N_{зв} \leq (2n - 3)$. Розв'язавши відповідні квадратичні нерівності відносно n , отримуємо діапазон $3 \leq n \leq 14$. Тобто, якщо кількість порівнюваних об'єктів менше 14, то у методі черг можна генерувати ПП строго у порядку черг, як показано у [3]. Якщо об'єктів більше 14, то у методі черг слід починати генерувати ПП зі структури, якій відповідає дводольний граф покривного дерева, де a_1 (найкращий об'єкт) порівнюється з об'єктами з 2-ї половини ранжирування $(a_{\lfloor \frac{n}{2} \rfloor + 1}, \dots, a_n)$, а a_n (найгірший об'єкт) – з об'єктами з 1-ї половини ранжирування $(a_1, \dots, a_{\lfloor \frac{n}{2} \rfloor})$. До цієї множини з $(n - 1)$ ПП слід додавати ПП у порядку черг, доки кількість ПП не досягне $(2n - 3)$.

З урахуванням особливостей методів, що порівнюються, алгоритм експериментального дослідження методу ПП на стійкість включатиме наступні етапи:

1. Генерування множини нормованих відносних ваг об'єктів, що оцінюються $(w_1, \dots, w_n)$
2. Побудова матриці ПП (МПП) на основі згенерованих відносних ваг $\{a_{ij} = \frac{w_i}{w_j}; i, j = 1..n\}$

3. Побудова неповної МПП з $(2n - 3)$ незалежних елементів на основі повної, відповідно до заданого методу, що оцінюється за стійкістю
4. «Зашумлення» згенерованої МПП $\{a'_{ij} = a_{ij} \pm \varepsilon; i, j = 1..n\}$
5. Обчислення відносних ваг об'єктів $(w'_1, \dots, w'_n)$ заданим методом (власного вектору, найменших квадратів, геометричного середнього [6], покривних дерев[5], іншим) на основі зашумленої МПП.
6. Обчислення різниці (ординальної та/або кардинальної) між початковим та отриманим векторам відносних ваг Δ . У якості показника такої різниці може бути обраний один з численних показників, що використовуються у [6].

Слід зазначити, що в контексті планування подібного експерименту актуальною лишається проблема появи реверсу рангів у векторі відносних ваг об'єктів після зашумлення початкової МПП.

Результати дослідження [3] вказують на те, що ПП з черг з меншими номерами є більш інформативними та точними. Відтак, пропонується модифікувати формулу зашумлення елемента МПП (крок 4 викладеного алгоритму) з урахуванням номеру черги, до якої він належить. Номер черги визначається відстанню між об'єктами, що порівнюються, у ранжируванні. Відповідно, для урахування порядку (черговості) виконання ПП, пропонується застосовувати опуклу комбінацію двох компонентів зашумлення, де один з компонентів зменшується пропорційно до відстані між порівнюваними об'єктами у ранжируванні. Відповідні формули для адитивного та мультиплікативного зашумлення елементів МПП на 4-му кроці запропонованого алгоритму наведені нижче:

$$a'_{ij} = a_{ij} \pm \left(\alpha \varepsilon + (1 - \alpha) \frac{\varepsilon}{|i-j|} \right); i, j = 1..n; 0 < \alpha < 1; \quad (2)$$

$$a'_{ij} = a_{ij} (1 \pm (\alpha \delta + (1 - \alpha) \delta / |i - j|)); i, j = 1..n; 0 < \delta < 1; 0 < \alpha < 1 \quad (3)$$

Висновки

На поточному етапі дослідження підготовлено підґрунтя для експериментального порівняння кількох методів експертного оцінювання (best-worst, Top 2, метод черг) за стійкістю до помилок експертів. Визначено умови, за яких такий експеримент може бути проведений, а його результати можна вважати достовірними. Запропоновано модифікований алгоритм

експерименту з оцінки стійкості методу, у якому моделюється весь цикл експертизи без залучення реальних експертів.

Поточне дослідження та подальші експерименти дозволять поліпшити якість процесу прийняття рішень у слабо-структурованих предметних областях з урахуванням когнітивних та алгоритмічних аспектів експертного оцінювання.

1) Rezaei, J. (2015). Best-worst multi-criteria decision-making method. *Omega*, 53, 49–57. <https://doi.org/10.1016/j.omega.2014.11.009>.

2) Szádóczi, Z., Bozóki, S., Juhász, P. et al. (2023). Incomplete pairwise comparison matrices based on graphs with averaged degree approximately 3. *Ann Oper Res* 326, 783–807. <https://doi.org/10.1007/s10479-022-04819-9>.

3) Andriichuk, O., Kadenko, S., & Tsyganok, V. (2024). Significance of the order of pair-wise comparisons in Analytic Hierarchy Process: an experimental study. *Journal of Multi-Criteria Decision Analysis*, 31(3-4), e1830. <https://doi.org/10.1002/mcda.183>.

4) Andriichuk, O. & Kadenko, S. (2022). Ranking- dependent Decision Support Methods for Weakly- structured Domains. *CEUR. Workshop Proceedings*, Vol. 3503, 32-41.

5) Kadenko, S., Tsyganok, V., Szádóczi, Z., & Bozóki, S. (2021). An update on combinatorial method for aggregation of expert judgments in AHP. *Production*, 31: 1-17. doi: 10.1590/0103-6513.20210045.

6) Szádóczi, Z., Bozóki, S., & Tekile, H. (2022). Filling in pattern designs for incomplete pairwise comparison matrices: (quasi-) regular graphs with minimal diameter. *Omega*, 107, doi:10.1016/j.omega.2021.102557.

МЕТОД МЕТРИЗАЦІЇ ВАГОВИХ КОЕФІЦІЄНТІВ КРИТЕРІЇВ ПРИ ЛЕКСИКОГРАФІЧНОМУ УПОРЯДКУВАННІ АЛЬТЕРНАТИВ

Григорій Гнатіснко

*Київський національний університет імені Тараса Шевченка
g.gna5@ukr.net*

Розглядається задача багатоатрибутного вибору. Пропонується метод визначення кількісних значень вагових коефіцієнтів критеріїв при застосуванні лексикографічного критерію для упорядкування багатоатрибутної множини альтернатив. Запропоновано обґрунтування визначення вагових коефіцієнтів критеріїв

на основі апостеріорного аналізу їх значень. Наводиться обґрунтування застосування запропонованого методу.

Ключові слова: багатоатрибутний вибір, лексикографічний критерій, вагові коефіцієнти, рейтинг, точна апроксимація

Вступ

Задачі лексикографічного упорядкування альтернатив є поширеним напрямком наукових досліджень [1]. Вони успішно застосовуються на практиці у різноманітних постановках і різним класам таких задач вітчизняні та іноземні дослідники закономірно приділяють увагу [2, 3]. Створення інструментарію для розв’язання поширених прикладних задач є актуальною і часто невдячною справою. Завжди знайдуться опоненти, які поцікавляться – навіщо потрібен черговий метод при наявності кількох десятків традиційних, які добре працюють. Відтак треба доводити ефективність нового метода, пояснювати його позитивні риси, демонструвати перевагу над попередніми підходами тощо.

Постановка задачі та математична модель

Нехай маємо задачу багатоатрибутного вибору

$$x_i^j \in X, i \in I = \{1, \dots, k\}, j \in J = \{1, \dots, n\}, X \in R^k \quad (1)$$

$$F = (f_1, \dots, f_k), f_i = f_i(x) = x, i = 1, \dots, k. \quad (2)$$

Необхідно здійснити упорядкування альтернатив виду (1) за критеріями виду (2). Причому, кількість атрибутів альтернатив та кількість критеріїв задачі співпадають. І нехай особа, що приймає рішення, вирішила застосувати для вирішення цієї проблеми лексикографічний критерій.

Не зменшуючи загальності, будемо вважати, що критерії упорядковані за важливістю таким чином:

$$f_1 \succ f_2 \succ \dots \succ f_{k-1} \succ f_k. \quad (3)$$

Відповідно, вагові коефіцієнти, які відображують відносну важливість критеріїв у ранжуванні (3), відповідають таким нерівностям:

$$\rho_1 > \rho_2 > \dots > \rho_{k-1} > \rho_k. \quad (4)$$

Вектори значень критеріїв для різних альтернатив будемо позначати

$$x^j = (x_1^j, \dots, x_k^j), j \in J = \{1, \dots, n\}, x_i^j \in \{0, 1, \dots\} \text{ для } \forall i, j: i \in I, j \in J. \quad (5)$$

За результатами застосування лексикографічного критерію здійснюється упорядкування об'єктів множини (1) за важливістю. Тобто, щодо структури множини альтернатив (1) стає відомим порядок важливості альтернатив. Більше нічого про вектори виду (5) не відомо. Але у такій ситуації прийняття рішень існує додаткова інформація про структуру множини (1).

Не зменшуючи загальності, будемо вважати, що в результаті застосування лексикографічного критерію (2), структура якого відображена формулою (3), має місце такий порядок:

$$x^1 \succ x^2 \succ \dots \succ x^{n-1} \succ x^n$$

Задача полягає у визначенні кількісних значень вагових коефіцієнтів (4).

Евристика Е1. Для визначення узагальнених показників кожної альтернативи виду (5) будемо застосовувати лінійну згортку:

$$F^j = F(x^j) = \sum_{i=1}^k \rho_i x_i^j.$$

Евристика Е2. При застосуванні кількісних вагових коефіцієнтів критеріїв має зберігатися порядок, встановлений шляхом застосування лексикографічного критерію. Тобто, має виконуватися система нерівностей:

$$F^1 \geq F^2 \geq \dots \geq F^{n-1} \geq F^n$$

або

$$\sum_{i=1}^k \rho_i x_i^1 \geq \sum_{i=1}^k \rho_i x_i^2 \geq \dots \geq \sum_{i=1}^k \rho_i x_i^{n-1} \geq \sum_{i=1}^k \rho_i x_i^n.$$

Точна апроксимація

Означення 1. Особливими точками для точної апроксимації структури переваг на множині критеріїв (2) є точки, що породжуються такими ситуаціями:

$$\sum_{i=1}^k \rho_i x_i^{j-1} \geq \sum_{i=1}^k \rho_i x_i^j, j = 2, \dots, n,$$

але при цьому має місце нерівність: $x_i^{j-1} < x_i^j$ для деяких індексів i .

Означення 2. Точна апроксимація – це ситуація, коли в особливих точках породжуються рівності, для перетворення яких у відношення порядку до відповідних вагових коефіцієнтів яких слід відняти деяке достатньо мале число.

Евристика Е3. Для застосування точної апроксимації і дотримання вимог строгого упорядкування може бути, у разі необхідності, експертно задане деяке достатньо мале фіксоване відхилення від знайдених значень вагових коефіцієнтів:

$$\rho_i = \rho_i - \varepsilon, i \in \{1, \dots, k\}, \varepsilon > 0.$$

Евристика Е4. Значення поправки до відповідних вагових коефіцієнтів може бути обчислене, наприклад, як зворотна величина до найбільшого рейтингу альтернатив, визначеного за результатами розв'язання задачі лексикографічного упорядкування.

Обґрунтування застосування запропонованого методу

– метод визначення вагових коефіцієнтів критеріїв при лексикографічному упорядкуванні альтернатив може бути корисним у багатьох практичних задачах:

– метризація кваліметричної шкали – визначення кількісної відповідності порядковому відношенню переваги;

– визначенні компетентності експертів, якщо це є можливим, важливим та доцільним;

– порівняння реальної структури упорядкованої множини альтернатив з експертно заданою;

– аналіз структур різних множин альтернатив, зокрема, на предмет напруженості ситуації, справедливості, чутливості тощо.

1. Волошин О.Ф., Машенко С.О. Моделі та методи прийняття рішень: навч. посіб. для студ. вищ. навч. закл. / О.Ф. Волошин, С.О. Машенко. – 3-є вид., перероб. – К.: «Видавництво Людмила», 2018. – 292 с.

2. Семенова Н.В., Ломага М.М., Семенов В.В. Існування розв'язків та метод розв'язання лексикографічної задачі опуклої оптимізації з лінійними функціями критеріїв. Допов. Нац. акад. наук Укр. 2020. № 12. С. 19—27. <https://doi.org/10.15407/dopovidi2020.12.019>

3. Iskander, M. G. (2023). On modeling a lexicographic weighted maxmin–minmax approach for fuzzy linear goal programming. *Mathematical Modeling and Computing*, 10 (1), 186–194. <http://jnas.nbuiv.gov.ua/article/UJRN-0001413245>

ВИКОРИСТАННЯ КАТАЛОГІВ ДАНИХ В УПРАВЛІННІ БЕЗПЕКОЮ КОРПОРАТИВНИХ ІНФОРМАЦІЙНИХ РЕСУРСІВ

Ірина Храмова, ORCID: 009-008-0029-7795

Інститут проблем реєстрації інформації НАН України

Проаналізовані окремі аспекти сучасних інструментальних засобів для створення каталогів цифрових інформаційних ресурсів, що сприяють підвищенню безпеки корпоративних даних.

Ключові слова: цифровий інформаційний ресурс, безпека даних, каталог даних.

Вступ

Оскільки цифрові інформаційні ресурси дедалі стають основоположними для інформаційної взаємодії в складних організаційних структурах, панівними концепціями роботи з даними залишаються сховища та озера даних. Їх поєднує ідея централізованості управління. Всі дані стікаються в якесь центральне сховище, звідки їх можуть забирати для своїх цілей різні команди. Останнім часом набуває популярності також децентралізована архітектура даних, так звані сітки даних (data mesh). Коли багато джерел, багато споживачів і всі джерела даних досить різноманітні не тільки в плані структури даних, але і в тому, що кожне джерело має дані різних типів і різних локалізацій (хмари даних, тощо), політики управління безпекою даних в такому операційному середовищі стають чи не ключовими. Правильний і добре підтримуваний перелік активів даних міг би бути хорошою відправною точкою для управління безпекою. Однак він не може впоратися зі швидкістю, обсягом і складністю даних у сучасних корпоративних ландшафтах даних.

Каталог даних як засіб підтримки безпеки інформаційних ресурсів. Сучасні каталоги даних, хоча і не панацея, але можуть значно сприяти захисту та безпеці корпоративних даних, надаючи інструменти та можливості, які підвищують безпеку даних, управління та відповідність стандартам. Ось деякі ключові аспекти їхньої ролі та місця в підтримці безпеки даних.

Виявлення та класифікація даних. Допомогають виявляти та класифікувати активи даних, спрощуючи ідентифікацію

конфіденційної інформації та застосування відповідних заходів безпеки. Це гарантує, що лише авторизований персонал може отримати доступ до критично важливих даних, зменшуючи ризик витоку даних

Контроль доступу та керування дозволами. Забезпечують надійні механізми контролю доступу, такі як керування доступом на основі ролей (англ. Role Based Access Control, RBAC). Це допомагає запобігти несанкціонованому доступу та гарантує, що конфіденційні дані будуть доступні лише авторизованим користувачам.

Визначення походження даних. Надають можливості визначення походження даних, що дозволяє організаціям відстежувати дані до їх джерела та гарантувати їх цілісність, також це має вирішальне значення для оцінки ризиків безпеки.

Маскування та анонімізація даних. Підтримують методи маскування та анонімізації даних, що дозволяє користувачам працювати з реальними даними, зберігаючи безпеку конфіденційної інформації.

Аудит і моніторинг. Дозволяють проводити аудит і моніторинг доступу та використання даних, щоб виявляти потенційні загрози безпеці та реагувати на них. Це включає в себе також інтеграцію з інструментами безпеки, щоб забезпечити комплексний підхід до захисту даних.

Відповідність вимогам і керування даними. Оптимізують дотримання нормативних вимог, надаючи централізовану систему керування даними. Вони допомагають вести точні записи доступу до даних і їх використання, забезпечуючи дотримання галузевих та корпоративних стандартів і правил безпеки.

Класифікація даних. Ключову позицію серед інших аспектів займає класифікація даних, яка є критично важливим процесом в управлінні та захисті даних. Класифікація приносить численні переваги.

Покращення безпеки даних: визначає, які дані є найбільш конфіденційними та потребують найвищого рівня захисту. Це дозволяє їм застосовувати надійні заходи безпеки, такі як шифрування та контроль доступу, спеціально для даних, які їх найбільше потребують, зменшуючи ризик витоку даних і несанкціонованого доступу.

Посилення відповідності нормам і стандартам: визначає, які дані підпадають під дію нормативних обмежень, як потрібно

обробляти певні типи даних і переконатися, що вони відповідають правилам і стандартам безпеки.

Ефективне керування даними: допомагає в політиках збереження даних, гарантуючи, що дані не зберігаються довше, ніж необхідно. Також в разі видалення даних допомагає впевнитися, що не залишилось небажаного сліду видалених даних.

Управління та пом'якшення ризиків: визначає, які дані є найбільш цінними та піддаються ризику. Маючи ці знання, можна визначити пріоритети своїх ресурсів і зосередити стратегії зменшення ризиків на найважливіших даних, зменшуючи потенційний вплив інцидентів, пов'язаних із даними.

Підвищення усвідомлення та відповідальності: підвищує обізнаність споживачів про різні типи даних, які вони використовують, і протоколи безпеки доступу.

Можливості засобів класу Data Catalog для підтримки безпеки даних. Як приклад візьмемо деякі із найбільш популярних за станом на поточний рік сучасних інструментальних засобів класу Data Catalog, та перелічимо їх можливості з точки зору підтримки згаданих вище аспектів сприяння безпеці даних (Табл. 1).

Табл. 1. Підтримки безпеки даних засобами класу Data Catalog

№ п/п	Засіб	Наявність функцій, що сприяють підтримці безпеки даних
Засоби з відкритим кодом		
1.	Apache Atlas	інтеграція RBAC, LDAP, ручна та автоматична класифікація, візуалізація та інтеграція безпеки на основі хмари.
2.	Amundsen	інтеграція RBAC, LDAP, ручна та автоматизована класифікація, візуалізація та інтеграція безпеки на основі хмари.
3.	DataHub	RBAC, LDAP, ручна та автоматична класифікація, візуалізація та інтеграція безпеки на основі хмари.
4.	Marquez	інтеграція RBAC, LDAP, ручну та автоматизовану класифікацію, візуалізацію та інтеграцію безпеки на основі хмари.
5.	OpenMetadata	RBAC, LDAP, ручна та автоматична класифікація, візуалізація та інтеграція безпеки на основі хмари.

Пропріетарні засоби		
6.	Alation	RBAC, аналітика поведінки користувачів, автоматична класифікація, керування метаданими, відстеження походження за допомогою ШІ та інтеграція з технологіями керування захистом інформації (SIEM) та запобігання витоку конфіденційної інформації (DLP).
7.	Informatica	розширений RBAC, шифрування, класифікація на основі машинного навчання, автоматичне походження, візуалізація та інтеграція з SIEM та DLP.
8.	Collibra	розширений RBAC, шифрування, ручна та автоматична класифікація, візуалізація походження та інтеграція безпеки на основі хмари.
9.	Каталог знань IBM Watson	контроль доступу на основі на основі штучного інтелекту (ШІ), класифікація на основі ШІ, відстеження походження на основі ШІ та інтеграція безпеки на основі ШІ.
10.	Ataccama	RBAC, керування доступом до даних, вбудоване профілювання даних, класифікацію, автоматизоване походження та керування даними, відповідність вимогам безпеки.

Висновки

Каталоги даних із підтримкою визначених в роботі функціональних аспектів, не тільки дозволяють подолати різноманітність численних цифрових інформаційних ресурсів, але й посилити безпеку циркулюючих даних, захистити конфіденційні дані та забезпечити дотримання стандартів і корпоративних нормативів в сфері інформаційної безпеки.

УТОЧНЕННЯ ОЦІНОК ВАГОВИХ КОЕФІЦІЄНТІВ ДЛЯ МЕТОДУ ПАРНИХ ПОРІВНЯНЬ

В. Полуциганова, ORCID: 0000-0002-9729-5786,

С. Смирнов, ORCID: 0000-0003-4190-5204

КПІ ім. І. Сікорського

medvika@ukr.net, sergsmr@gmail.com

В роботі досліджується проблема уточнення експертних оцінок при визначенні вагових коефіцієнтів у методі парних порівнянь. Автори використовують метод множників Лагранжа для підвищення точності та об'єктивності отриманих результатів. За допомогою цього методу формалізовано процес уточнення вагових коефіцієнтів як математичну задачу оптимізації, що дозволяє врахувати додаткові обмеження, такі як нерівності між ваговими коефіцієнтами або залежності між критеріями.

Ключові слова: парні порівняння, вагові коефіцієнти, метод Лагранжа, оптимізація, уточнені оцінки, експертні методи

Вступ

Метод парних порівнянь (МПП) є потужним інструментом підтримки прийняття рішень [1]. Однак якість рішень значною мірою залежить від точності визначення вагових коефіцієнтів. Ці коефіцієнти відображають відносну важливість різних елементів, що порівнюються та обчислюються на основі суджень експертів. У цій роботі ми розглянемо, як можна уточнити вагові коефіцієнти за допомогою методу множників Лагранжа, що дозволяє враховувати існуючі обмеження та нормування.

Проблема уточнення оцінок вагових коефіцієнтів

Основним джерелом неточності при визначенні вагових коефіцієнтів є суб'єктивні помилки експертних оцінок. Різні експерти можуть піддаватися впливу різних факторів та мати різні погляди на відносну важливість елементів порівняння, що призводить до помилок та розбіжностей в оцінках [2]. Крім того, матриці парних порівнянь можуть містити неузгодженості внаслідок неувважності та втоми, що також спотворює результати.

Вимагати від експертів точності оцінок при проведенні великої кількості порівнянь є занадто оптимістично та зовсім нереалістично.

Метод Лагранжа для уточнення оцінок

Метод множників Лагранжа є потужним математичним інструментом для вирішення задач оптимізації з обмеженнями. У контексті МПП його можна застосувати для уточнення вагових коефіцієнтів наступним чином [3]:

1. Формулювання критерію оптимальності та обмежень;
2. Побудова функції Лагранжа;
3. Пошук її екстремуму, розв'язання системи рівнянь.

До переваг використання методу Лагранжа входить можливість врахування додаткових умов щодо значень вагових коефіцієнтів, оптимізація рішення відповідно до необхідного критерію, зручність та математична строгість, за наявності аналітичного розв'язку задачі.

Уточнення оцінок вагових коефіцієнтів для методу парних порівнянь забезпечена тим, що на відміну від класичного підходу МПП експертам пропонується виконати мінімально необхідний набір порівнянь (n замість $n(n-1)/2$). Експерти не перевантажуються та менше помиляються, і це дозволяє отримати точніші оцінки. Нехай w_i шукані вагові коефіцієнти, α_{ij} – надані експертні оцінки відношення ваг $w_i/w_j \approx \alpha_{ij}$, але внаслідок помилок $\prod \alpha_{ij} \neq 1$, хоча точне відношення ваг гарантує добуток значення 1. Позначимо β_{ij} відношення оцінок ваг та фіксуємо обмеження [4]:

$$\begin{cases} \sum_{i=1}^n w_i = 1 \\ \prod \beta_{ij} = 1 \\ \frac{w_i}{w_j} = \beta_{ij} \end{cases}$$

Для оцінки вагових коефіцієнтів використаємо метод найменших квадратів:

$$\begin{cases} \sum_{i=1}^n w_i = 1 \\ 0 < w_i < 1 \\ \sum (\frac{w_i}{w_j} - \alpha_{ij})^2 \rightarrow \min \end{cases}$$

Для вирішення цієї задачі використовуємо метод Лагранжа та отримуємо загальний вигляд рішення, для спрощення запису позначимо $\beta_{ij}(\beta_{ij} - \alpha_{ij}) = \mu$

$$\beta_{ij}^{1,2} = \frac{1}{2}(\alpha_{ij} \pm \sqrt{\alpha_{ij}^2 + 4\mu})$$

Розглядаються два можливі випадки:

$$\text{А) } \prod \alpha_{ij} > 1, \beta_{ij} < \alpha_{ij} \Rightarrow \mu < 0 ;$$

$$\text{Б) } \prod \alpha_{ij} < 1, \beta_{ij} > \alpha_{ij} \Rightarrow \mu > 0 .$$

В обох випадках з коренів квадратного рівняння обираємо той що з плюсом перед детермінантом, адже він очевидно ближчий до α_{ij} аніж корінь з мінусом, бо за допомогою МНК ми намагаємося мінімально відкоригувати наявні експертні оцінки.

Тоді:

$$\prod \beta_{ij} = 1 \Rightarrow \prod (\alpha_{ij} + \sqrt{\alpha_{ij}^2 + 4\mu}) = 2^n$$

де n – кількість множників.

Отже, при відомих α_{ij} , рівняння для β_{ij} хоч і є ірраціональними, але з однією змінною, розв'язуються чисельно та надають більш коректні оцінки вагових коефіцієнтів.

Висновки

Пропонований метод дозволяє уточнити оцінки вагових коефіцієнтів для МПП. За допомогою методу множників Лагранжа розроблено ефективну процедуру обчислення оцінок ваг на основі мінімального набору парних порівнянь. Він дозволяє враховувати наявні обмеження, забезпечує точність та самоузгодженість оцінок.

1. R. W. Saaty, “The analytic hierarchy process—what it is and how it is used”, *Math. Modelling*, т. 9, № 3-5, с. 161–176, 1987. Дата звернення: 27 листоп. 2024. [Онлайн]. Доступно: [https://doi.org/10.1016/0270-0255\(87\)90473-8](https://doi.org/10.1016/0270-0255(87)90473-8)

2. D. Jato-Espino, E. Castillo-Lopez, J. Rodriguez-Hernandez та J. C. Canteras-Jordana, “A review of application of multi-criteria decision making methods in construction”, *Automat. Construction*, т. 45, с. 151–162, верес. 2014. Дата звернення: 27 листоп. 2024. [Онлайн]. Доступно: <https://doi.org/10.1016/j.autcon.2014.05.013>

3. M. S. Bazaraa, H. D. Sherali та C. M. Shetty, *Nonlinear Programming*. Hoboken, NJ, USA: Wiley, 2005. Дата звернення: 27 листоп. 2024. [Онлайн]. Доступно: <https://doi.org/10.1002/0471787779>

4. В. І. Полуциганова та С. А. Смирнов, “Робасний метод оцінок вагових коефіцієнтів для MAUT”, у *Інформ. технології та безпека*, Київ, Україна, 30 листоп. 2023. Київ: ІН-Т ПРОБЛЕМ РЕЄСТРАЦІЇ ІНФОРМАЦІЇ НАН УКРАЇНИ, 2023, с. 50–54.

ВИЗНАЧЕННЯ ВЗАЄМНОЇ ПОЗИЦІЇ ДРОНІВ ЗА УМОВ ВІДСУТНОСТІ GPS-СИГНАЛУ ДЛЯ ПІДТРИМКИ ПРИЙНЯТТЯ РІШЕНЬ ПІД ЧАС ВИКОНАННЯ ПОСТАВЛЕНОЇ ЗАДАЧІ

Вячеслав Корольов, ORCID: 0000-0003-1143-5846,
korolev.academ@gmail.com,

Олег Рибальченко, ORCID: 0000-0002-5716-030X,
rv.oleg.ua@gmail.com,

Олександр Ходзінський, ORCID: 0000-0003-4574-3628,
okhodz@gmail.com,

Максим Огурцов, ORCID: 0000-0002-6167-5111,
maksymogurtsov@gmail.com,

Олександр Ярушевський, ORCID: 0009-0004-9930-1643,
olexand777@gmail.com,

Інститут кібернетики імені В.М. Глушкова НАН України

Останнім часом групи дронів значно розширили сферу застосування, включаючи такі завдання, як зрошення полів, виявлення шкідників, доставлення вантажів та моніторинг стану рослин. Для підвищення оперативної ефективності та швидкості на великих територіях все частіше використовують групи дронів. Однак у віддалених регіонах, де сигнали глобальної системи позиціонування (GPS) можуть бути неточними або недоступними, виникає потреба в ефективних методах визначення точних відносних положень дронів у групі.

Метою цієї роботи є розробка алгоритмів відносного позиціонування для груп дронів в умовах обмеженого доступу до GPS-сигналів. Ці алгоритми повинні дозволити визначення відносних положень дронів та загальної

конфігурації групи навіть при повній відсутності GPS-сигналів.

В результаті дослідження був розроблений алгоритм визначення відносного положення елементів рою в умовах відсутності глобальних сигналів позиціонування та відповідний математичний апарат. Він враховує можливу неточність вимірювання відстані.

Отримані результати були успішно підтверджені практичними розрахунками.

Ключові слова: дрони, БпЛА, GPS, позиціонування, підтримка прийняття рішень.

Вступ

За останні роки дрони значно розширили сферу застосування, включаючи такі завдання, як зрошення полів, виявлення шкідників, доставлення вантажів та моніторинг стану рослин [1]. Для підвищення ефективності та швидкості операцій на великих площах все частіше використовують групи дронів. Однак на територіях, де сигнали глобальної системи позиціонування (GPS) можуть бути неточними або недоступними, виникає потреба в ефективних методах визначення точних відносних положень дронів у групі, щоб запобігти зіткненням і забезпечити можливість виконання завдань [2].

Мета

Метою роботи є розробка алгоритмів визначення взаємної позиції для груп дронів в умовах обмеженого доступу до GPS-сигналів. Ці алгоритми повинні дозволити визначення відносних положень дронів і загальної конфігурації групи навіть при повній відсутності GPS-сигналів або при наявності GPS-сигналів лише у оператора. Алгоритми повинні працювати навіть у випадку великих відстаней між дронами.

Аналіз попередніх досліджень

Рішення, які використовуються на сьогоднішній день для вирішення цієї проблеми, зазвичай покладаються на дороги мережі стаціонарних станцій (якорів) з точно відомими позиціями [3].

Ця проблема є постійним предметом наукових досліджень, про що свідчить велика кількість опублікованих робіт на цю

тому. Але більшість з них призначена для іншого класу задач – точного позиціонування малих дронів всередині приміщення.

Цікавим підходом є запропонований в роботі [4] – пропонується використання сірої низькороздільної камери та технології комп'ютерного зору, щоб використовувати її з нейронною мережею, однак цей підхід вимагає чималих обчислювальних ресурсів.

У нашому попередньому дослідженні [5] використовувалися гранулярні обчислення та нечіткі множини, але вони не забезпечили достатньої точності для застосування до поточної задачі.

Як ми бачимо, існуючі рішення не змогли фактично вирішити проблему, враховуючи всі її вимоги та обмеження. Тому розглянемо, як можна визначити відносні положення дронів і загальну конфігурацію групи навіть за умови повної відсутності GPS-сигналів.

Фундаментальний підхід, який пропонується використати – визначення напрямку та відстані від кожного дрона до кожного іншого дрона.

Основна частина

Для визначення відстані між парою дронів пропонується використовувати або вимірювання сигналу RSSI, або вбудовані в дрон інструменти вимірювання відстані, які доступні на багатьох моделях. Однак у будь-якому випадку необхідно враховувати наявну неточність цього вимірювання відстані.

А для точного визначення кутового положення кожного дрона відносно іншого недостатньо вимірювати лише відстань між ними. Пропонується розгорнути два дрони, що дозволить розрахувати їх відповідні кутові положення.

Розглянемо вимірювання відстаней від дронів А і В до дрону Т (ціль), відстань і кут до якого необхідно визначити. Для цього малюються два уявні кола з центрами в точках А і два кола з центрами в точці В. Ці кола представляють відстані до Т плюс або мінус похибку вимірювання. Таким чином дрон Т не розташований у конкретній точці, а в певній зоні, утвореній перетином цих кіл. Відстань між дронами А і В визначається за допомогою тих самих методів, що й відстань до дрона Т.

Використовуючи формули тригонометрії (зокрема, теорему Піфагора і теорему косинусів), ми змогли визначити позицію дрона Т відносно А і В з великою точністю.

Тестування алгоритму

Отримані формули були експериментально перевірені шляхом створення кількох тестових наборів даних для різних відносних положень трьох дронів та різних відстаней між ними. Результати практичного тестування підтвердили працездатність створеної математичної моделі, дозволяючи визначити відносне положення дронів у групі відносно один одного.

Висновки

Автономні елементи групи дронів дозволяють здійснювати точний моніторинг та управління завданнями групового використання дронів на великих площах, оптимізуючи використання ресурсів та зменшуючи трудові витрати. В результаті цього дослідження був розроблений алгоритм визначення відносного положення елементів групи дронів в умовах відсутності GPS-сигналів та відповідний математичний апарат. Він враховує можливу неточність вимірювання відстані.

Результати були перевірені за допомогою практичних обчислень. Тестування підтвердили працездатність створеної математичної моделі, забезпечуючи похибку вимірювання 1% або менше.

Майбутні дослідження будуть зосереджені на практичному застосуванні та емпіричній оцінці отриманих результатів в групах дронів.

[1] M. R. Dileep, A. V. Navaneeth, S. Ullagaddi, and A. Danti, A study and analysis on various types of agricultural drones and its applications, in: Proceedings of the 2020 fifth international conference on research in computational intelligence and communication networks (ICRCICN), IEEE, 2020, pp. 181–185. doi:10.1109/ICRCICN50933.2020.9296195

[2] D. Albani, J. Ijsselmuiden, R. Haken, and V. Trianni, Monitoring and mapping with robot swarms for agricultural applications, in: Proceedings of the 2017 14th IEEE International Conference on Advanced Video and Signal Based Surveillance (AVSS), IEEE, 2017, pp. 1-6. doi:10.1109/AVSS.2017.8078478

[3] P. Perazzo, F. B. Sorbelli, M. Conti, G. Dini, C. M. Pinotti. Drone path planning for secure positioning and secure position verification. IEEE Transactions on Mobile Computing 16, no. 9 (2016): 2478-2493. doi:10.1109/TMC.2016.2627552

[4] Crupi, Luca, Alessandro Giusti, and Daniele Palossi. "High-throughput visual nano-drone to nano-drone relative localization using

onboard fully convolutional networks." arXiv preprint arXiv:2402.13756 (2024).

[5] M. Ogurtsov, V. Korolyov, O. Khodzinskyi, Improving the Productivity of UAV Operations Based on Granular Computing and Fuzzy Sets, in: Proceedings of the 2021 IEEE 6th International Conference on Actual Problems of Unmanned Aerial Vehicles Development (APUAVD) 19-21 Oct. 2021, 2021, pp. 33-36. doi:10.1109/APUAVD53804.2021.9615419.

КЛАСИФІКАЦІЯ ОСНОВНИХ ЗАДАЧ ВЕБ-ЗАСТОСУНКІВ

Олексій Веретьонкін, ORCID: 0009-0007-4776-5333

*Київський національний університет імені Тараса Шевченка
alexveretenkin@gmail.com*

Метою доповіді є систематизація та аналіз основних задач веб-застосунків шляхом їх класифікації, а також визначення ключових аспектів для кожної категорії задач. Представлена класифікація дозволяє систематизувати процес аналізу та розробки, що сприяє підвищенню ефективності та якості кінцевого продукту.

У контексті аналізу веб-застосунків основна проблема полягає в тому, що без чіткої класифікації задач важко проводити ефективний аналіз та розробку. Це стосується як планування ресурсів, так і вибору технологій та архітектурних рішень. Формулювання проблеми в тезах підкреслює необхідність систематизації знань та підходів для подолання цих викликів.

Результати дослідження полягають у наступній класифікації (список задач може розширюватися, тож було обрано лише основні задачі):

- Задачі що реалізують бізнес-логіку
 - Авторизація та автентифікація
 - Обробка даних
 - Інтеграція з зовнішніми сервісами
 - Контент-менеджмент
 - Спеціальне відображення даних (графіки, діаграми, 3d-моделі тощо)
- Нефункціональні задачі або задачі що не відносяться до бізнес-логіки
 - Масштабованість
 - Підтримка та обслуговування

- Продуктивність
- Безпека
- Зручність використання системи (UI/UX)

Ідея такої простої класифікації полягає в тому, що вона дозволяє швидко визначити з яким типом веб-застосунку ми маємо справу - веб-додаток що являє собою сайт (web-site), чи веб-застосунок як додаток (web-app) [1].

Запропонована класифікація задач веб-застосунків дозволяє систематизувати процес аналізу та розробки, що підвищує ефективність та якість кінцевого продукту. Чітке розуміння категорій задач сприяє оптимальному вибору технологій, покращенню безпеки, продуктивності та користувацького досвіду.

1. Geers M. Micro Frontends in Action / Michael Geers., 2020. – 296 с. – (Manning).

МЕРЕЖЕВИЙ ВПЛИВ ДАНИХ НА КОРПОРАТИВНІ ЦІЛІ

Олексій Гайворонський¹, ORCID: 0000-0002-8470-8780,
Alexei.Gaivoronski@ntnu.no

Василь Горбачук², ORCID: 0000-0001-5619-6979,
GorbachukVasyl@netscape.net

Максим Дунаєвський², ORCID: 0000-0002-6926-398X,
MaxDunaievskiy@gmail.com

Сеїт-Бекір Сулейманов², ORCID: 0000-0002-7141-6113,
¹*Norwegian University of Science and Technology,
Trondheim, Norway*

²*Інститут кібернетики імені В.М.Глушкова НАН України, Київ,
Україна, SBSuleimanov@gmail.com*

Щоб підтримувати і зберігати конкурентну перевагу, промислові підприємства мають неперервно адаптувати свою бізнес-модель до змінюваних вимог ринку, регуляторних положень і поступу технологічних інновацій.

Ключові слова: менеджмент ресурсу даних, керовані даними продукти і послуги, модель.

Вступ

Вимоги ринку матеріалізуються головним чином у запитах клієнтів, які постійно розвиваються. Крім ціни та якості, стають важливими екологічні вимоги.

Кількість і сфера регуляторних положень збільшується через зростаючу важливість екологічних і суспільних умов, в яких працює промислове підприємство. З 2023 р. набув чинності Закон про ланцюги постачання (Supply Chain Act) в умовах європейського регулювання економіки даних на Єдиному (цифровому) ринку (наприклад, Закон про врядування даних (Data Governance Act) і Закон про дані (Data Act)). Європейська стратегія даних (European Data Strategy) спирається на регулювання даних в Європі та створення просторів даних.

Технологічні інновації є одним із найпотужніших джерел конкурентної переваги, зокрема для підприємств у високорозвинених країнах з високою вартістю робочої сили, в яких конкуренція за витратами зазвичай не є можливою. Технологічні інновації матеріалізуються в різних формах і передбачають інновації товарів і послуг, а також інновації, породжені даними. Рентабельність (profit margin) підприємств послуг перевершує норму прибутку товарно-центричних підприємств, а технологічна компанія вважається перспективною для ринкової капіталізації (Tesla, Zalando та інші). Сьогодні ІІІ та дані є ресурсами цифрових послуг, які генерують і генеруватимуть найвищі темпи зростання. Тому у менеджменті досліджувалося і досліджується питання, як промислове підприємство має реагувати на технологічні зміни.

Основна частина

Виділяється п'ять сил, що визначають корпоративну стратегію [1], і проводиться відповідний аналіз конкурентних переваг [2]. Динамічні спроможності дозволяють промисловим підприємствам неперервно реорганізовувати свою ресурсну базу для підтримки конкурентної переваги. Сьогодні такі спроможності вважаються теоретичною структурою для пояснення успішної (цілеспрямованої) адаптації ресурсної бази організації (підприємства) до змінюваного середовища [3]. Динамічні спроможності – це особлива форма організаційних спроможностей, які загалом дозволяють користуватися ресурсною базою організації. Так званий оснований на ресурсах

погляд на підприємство є теоретичною структурою у менеджменті, яка використовується для визначення стратегічних ресурсів, які фірма може експлуатувати для досягнення самопідтримуваної конкурентної переваги [4].

З цього погляду дані вважаються стратегічним ресурсом для конкурентної переваги промислового підприємства.

Практики менеджменту мають застосовуватися до ресурсів підприємства, щоб діставати, використовувати, розташовувати, відновлювати їх оптимальним способом, вимірюваним відповідно до вартості, часу та якості. Важливо розуміти, що дані типово не дають безпосереднього ділового виграшу, але уможливають ділові виграші через мережу механізмів впливу (effect mechanisms) [5].

Можна показати графічно мережу залежностей виграшів (Benefit Dependency Network) для каталогів даних, які типово запроваджуються для обробки метаданих як єдиного джерела істини (Single Source of Truth, SSOT) і для підвищення якості даних поміж підрозділів підприємства. Така мережа як методологія менеджменту виграшів була розроблена у 1990-х роках в Кренфілдському університеті Сполученого Королівства. Ця мережа містить 5 рівнів:

- 1) засоби реалізації (enablers);
- 2) уможливлені зміни (enabled changes);
- 3) ділові та індивідуальні зміни;
- 4) виграші;
- 5) інвестиційні цілі.

Кожний рівень містить 3 підрівні:

- 1.1) настанови обміну даними підприємства;
- 1.2) каталог даних;
- 1.3) EDM-фреймворк;
- 2.1) довіра до даних;
- 2.2) збільшена доступність даних завдяки інтеграції даних;
- 2.3) зручна інтеграція метаданих в 1.2);
- 3.1) чіткі відповідальності (зокрема через врядування даних);
- 3.2) збільшена грамотність з даних серед працівників;
- 3.3) автоматизовані процеси для створення записів даних;
- 4.1) поліпшення процесів;
- 4.2) повне і докладне подання E2E-процесу;
- 4.3) підвищена якість даних;
- 5.1) операційна досконалість;

5.2) дотримання регулювань (зокрема Закону ЄС про дані, Закону ЄС про ланцюги постачання тощо);

5.3) сприяння вдоволенню (contentment) клієнтів (B2B, B2C).

Мережу (орієнтований граф) формують орієнтовані ланки:

1.2)→2.1); 1.2)→2.2);
2.1)→3.1); 2.1)→3.2);
2.2)→3.3);
2.3)→3.3);
3.1)→4.1); 3.1)→4.2); 3.1)→4.3);
3.2)→4.1); 3.2)→4.3);
3.3)→4.2); 3.3)→4.3);
4.1)→5.1); 4.1)→5.2);
4.2)→5.1); 4.2)→5.2); 4.2)→5.3);
4.3)→5.1); 4.3)→5.2); 4.3)→5.3).

Діаграма перегляду показує, що запровадження каталогів даних не має прямого впливу на інвестиційні цілі, але натомість веде до організаційних і ділових змін, які, в свою чергу, ведуть до ділових вигравів, зокрема надійності. В цьому полягає менеджмент даними підприємства.

Висновки

Немає монопричинного взаємозв'язку між діяльностями та інструментами менеджменту ресурсами даних, з одного боку, та бажаними стратегічними цілями підприємства, з іншого боку, зокрема цілями безпеки. Підприємства мають усвідомлювати, що вплив такого менеджменту стає дієвим через складну мережу причинно-наслідкових механізмів. Оскільки дані є не матеріальним (tangible), а нематеріальним (immaterial) ресурсом, то поводження з даними відрізняється від поводження з матеріалами.

1. Porter M.E. *Competitive Strategy*. Free Press, 1986.
2. Porter M. E. *Competitive Advantage*. Free Press, 1998.
3. Teece D. J., Pisano G., Shuen A. Dynamic capabilities and strategic management. *Strategic Management Journal*. 1997. 18 (7). P. 509–533.
4. Гайворонський О.О., Горбачук В.М., Дунаєвський М.С. Стратегічна взаємодія провайдерів диференційованих Інтернет-послуг. *Проблеми управління та інформатики*. 2021. № 6. С. 102–113.
5. Gaivoronski A., Gorbachuk V., Dunaievskyi M., Suleimanov S.-B. Digital platforms to close the information asymmetry gaps. *Problems of Control and Informatics*. 2022. № 6. P. 67–82.

РОЗУМНИЙ ЕНЕРГОЗБЕРІГАЮЧИЙ ПРОМИСЛОВИЙ ТЕПЛИЧНИЙ КОМПЛЕКС З ВИКОРИСТАННЯМ ТЕХНОЛОГІЇ ІОТ НА БАЗІ CISCO PACKET TRACER

Микола Кіктев¹ [0000-0001-7682-280X], Роман Поліщук¹ [0009-0009-0356-0226],
Сергій Шворов¹ [0000-0003-3358-1297], Вячеслав Івашук¹ [0000-0002-9678-0990],
Павло Мазурчук¹ [0009-0000-0064-8999]

¹ Національний університет біоресурсів і природокористування
України, Київ, Україна
nkiktev@ukr.net, polishchuk.r23@gmail.com, sosdoc@i.ua,
v.ivaschuk@nubip.edu.ua

У докладі розглянуто створення розумних енергозберігаючих теплиць з використанням локальних ресурсів, таких як біогазові реактори та централізована енергомережа країни. Для інтелектуального управління енергопостачанням теплиці використовуються різні датчики та веб-камери. Програмне забезпечення Cisco packet tracer використовується для представлення «розумного тепличного комплексу» з використанням технології IoT і дозволяє моделювати реальну мережу, включаючи компоненти, пристрої та протоколи. Інтерфейс дозволяє будувати топологію та фізично розміщувати пристрої під контролем за допомогою смартфона та ноутбука. Управління системою організовано за допомогою спеціального модуля, адаптованого до рішень IoT, який може отримувати інформацію не лише від датчиків, а й від веб-сервісів. Модуль керування дозволяє реалізовувати складні сценарії, що критично важливо, враховуючи біологічну складову об'єктів господарювання теплиць.

Ключові слова: Інтернет речей, теплиця, енергозбереження, модуль керування, сервер, мережа, топологія.

Вступ

Виробництво біогазу є актуальним трендом для зеленої енергетики та сільськогосподарського виробництва, в тому числі для тепличних господарств, де можна досягти максимального

економічного ефекту від виробництва. Це пов'язано з тим, що технології закритого ґрунту потребують як електроенергії для освітлення, тепла для обігріву, так і CO_2 для живлення рослин, який утворюється при спалюванні газу. Проводяться дослідження щодо можливості використання як добрив твердих і рідких фракцій продуктів, які утворюються на біогазових установках (БГУ) під час генерації біогазу [1]. Специфікою виробництва біогазу на території України є тривалий зимовий період, коли необхідно витратити енергію на підтримку процесу газоутворення [2]. Авторами запропоновано запровадити ринкові умови щодо вибору залучення зовнішніх та внутрішніх джерел електроенергії [3]. Залежність вартості комерційної електроенергії в період 1-7 лютого 2024 р. є за посиланням: <https://www.oree.com.ua/index.php/pricestr>. З наведених даних видно, що тричі протягом тижня вартість електроенергії на кілька годин була низькою і було доцільно використовувати її для обслуговування БГУ. Зміна вартості електроенергії є випадковою, тому для прийняття рішення про залучення комерційної енергії необхідно враховувати багато показників: температуру в БГУ, температуру сировини перед завантаженням та ін. Збір інформації на сучасному технічному рівні можливий завдяки впровадженню технологій IoT з інтеграцією до технологій BigData для управління складними бізнес-процесами.

Метою дослідження – розробка інтелектуальної системи енергопостачання промислової теплиці із спільним використанням централізованих джерел енергопостачання та локальних джерел, реалізованих на базі біогазової установки.

Методологія (опис інтелектуальної системи)

Методом реалізації інтелектуальної системи є програмування та налаштування ключових компонентів, що забезпечують її працездатність. Для цього було використано два основні методи: налаштування через графічний інтерфейс користувача (GUI) на вкладці Config та через інтерфейс командного рядка (CLI). Ці підходи дозволяють ефективно керувати маршрутизаторами, шлюзами та сервером IoT, які є важливими елементами забезпечення стабільного та безперебійного енергопостачання теплиці. Шлюз і IoT-сервер налаштовано на взаємодію зі службою аутентифікації, що дозволяє не тільки відстежувати стан пристроїв, підключених до IoT-сервера через Інтернет, але й реалізовувати складні правила та умови для автоматизованого

керування смарт-об'єктами в теплиці. Це включає логіку програмування для реагування на певні зміни температури або вологості, які можуть вимагати автоматизованого втручання, наприклад, активації системи жалюзі або подачі біогазу для виробництва тепла та електроенергії.

Результати досліджень

Продемонструємо графічний інтерфейс системи управління та збору інформації, який ефективно відображає всі компоненти першого блоку, включаючи датчики температури та системи штор (рис. 1, а). Інтерфейс одного з термостатів ілюструє деталі програмування, залучених до його роботи. Також зображено систему жалюзі, інтегровану в автоматизовану систему керування з можливістю ручного налаштування через інтерфейс користувача. Ця можливість ручного регулювання надає системі важливу гнучкість, дозволяючи оператору швидко адаптуватися до мінливих умов вирощування, що є критично важливим для забезпечення оптимального мікроклімату в теплиці та ефективного управління ресурсами. На рис. 1, б показано імітацію роботи системи яка мала переключитися з споживання газу від біогазової установки при умові низької вартості енергії ($\leq 3\,000$ грн/МВт) від центральної мережі електропостачання.

Висновки

У цьому дослідженні було реалізовано симуляційний прототип «розумного парника» на основі Інтернету речей, у якому різноманітні сервіси відстежуються та контролюються через графічний інтерфейс веб-сторінки на смартфонах, ноутбуках і ПК. Для реалізації «розумної теплиці» використовуються принципи гнучкості та масштабованості, оскільки можна змінювати пристрої та сервіси, додавати або видаляти нові системи. Використовується програмне забезпечення відстеження пакетів Cisco, оскільки воно включає роботу з IoT і за допомогою якого можна моделювати реальну мережу, включаючи компоненти, пристрої та протоколи. Робочий простір у системі відстеження пакетів дозволяє будувати топологію та фізично розташовувати пристрої.

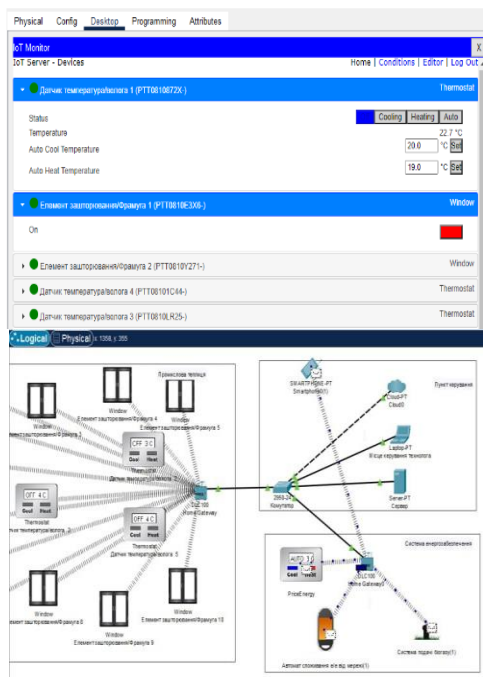


Рис. 1. Сполучені пристрої умовної моделі промислової теплиці (а); інтерфейси Cisco packet tracer в режимі симуляції

1. D. A. Meza-Carhuamaca, M. V. Gavilan-Quispe, E. F. Zevallos-Jara and E. N. Gabriel-Campos, "Automated Biogas Production for Electricity Supply at CEPAP," *2023 6th International Conference on Robotics, Control and Automation Engineering (RCAE)*, Suzhou, China, 2023, pp. 470-476, doi: 10.1109/RCAE59706.2023.10398829.

2. M. Spodoba and O. Spodoba, "Mathematical Model of Changes in Energy Costs for Thermostabilization of the Substrate and Objects in a Biogas Reactor," *2023 IEEE 5th International Conference on Modern Electrical and Energy System (MEES)*, Kremenchuk, Ukraine, 2023, pp. 1-6, doi: 10.1109/MEES61502.2023.10402431.

3. S. A. Shvovor, N. A. Pasichnyk, O. A. Opryshko, D. S. Komarchuk, A. O. Dudnyk and F. V. Hluhan, "The Methodological Foundations of Building an Energy Efficient Community," *2022 IEEE 16th International Conference on Advanced Trends in Radioelectronics, Telecommunications and Computer Engineering (TCSET)*, Lviv-Slavske, Ukraine, 2022, pp. 297-300, doi: 10.1109/TCSET55632.2022.9766956.

ЗМІСТ

<i>Олександр Додонов, Олена Горбачик, Марина Кузнєцова</i> Способи і заходи забезпечення функціональної стійкості інформаційних систем критичних інфраструктур	3
<i>Вячеслав Петров, Андрій Крючин, Цзичунь Ле, Євген Беляк</i> Організація оптичних систем багаторівневого об'ємного запису для довгострокового збереження великих масивів даних	7
<i>Дмитро Ланде, Леонард Страшной</i> Black Hat AI - виклики та шляхи протидії	10
<i>Oleksii Novikov, Iryna Stopochkina, Andrii Voitsekhovskiy, Mykola Ilin</i> Attack Models for Industrial Control System Elements Based on a Graph-Based Approach and Countermeasures	16
<i>Valerii Sytnikov, Dmytro Kalynenko</i> High-Performance AI Approaches for Real-Time Video Quality Improvement	20
<i>Houda El Bouhissi, Tetiana Hovorushchenko, Olga Pavlova</i> Hate Speech Detection: a hybrid approach	23
<i>Volodymyr Zhydkov</i> Finding defects in periodic structures	38
<i>Olexandr Polishchuk, Dmytro Polishchuk</i> Protection of multilayer network systems from targeted group attacks on the process of intersystem interactions	41
<i>Андрій Болдак, Костянтин Єфремов, Олександр Стусь, Олег Дмитренко</i> Використання семантичних мереж для виявлення потенційних фейків	45
<i>Юлія Рогушина, Анатолій Гладун, Сергій Прийма, Олена Аніщенко</i> Побудова профіля андрагога на основі відкритих джерел інформації	49
<i>Євген Лабун, Оксана Золотухіна</i> Створення датасету для навчання моделі нейронної мережі для мастерингу аудіо	52
<i>Володимир Юзефович, Євгенія Цибульська, Ніколай Стоянов</i> Ідентифікація інформаційних об'єктів різних джерел при формуванні інформаційного ресурсу системи моніторингу динамічних об'єктів навколишнього простору	55

Дмитро Манько, Євген Беляк

**Особливості організації крос-платформної системи
розпізнавання візуальних об'єктів у режимі реального часу** 59

Микола Кошевий, Павло Ломовцев

**Геометричне моделювання набору мобів для комп'ютерної
гри у жанрі sci-fi** 62

Олексій Безуглий

**Розробка алгоритму кластеризації об'єктів у QGIS з
використанням просторових даних** 64

Юрій Даник

**Особливості взаємозалежності та ризикології кіберпростору
і штучного інтелекту** 67

Анна Гончарова, Оксана Золотухіна

**Аналіз методів реферування новинних статей на основі
метаданих за допомогою штучного інтелекту** 71

Ірина Гришанова, Юлія Рогушина

**Планування послідовностей сервісів навчальних об'єктів
на основі машинного навчання з підкріпленням** 74

Володимир Мальчиков

**Метод вибору нескоротного коефіцієнта масштабування
недіадного вейвлет-перетворення** 78

Микола Кіктєв, Денис Жук, Олександр Ромащук,

Дудник Алла, Олексій Опришко

**Методика попередньої обробки зображення для
ідентифікації воронок від вибухів на полях** 81

Олег Гуменюк, Іван Золотов

**Графова модель аудиту кібербезпеки з використанням
генеративного штучного інтелекту** 84

Ігор Данилюк, Володимир Куцаєв

**Дослідження можливостей інтеграції квантових технологій
у наявні військові телекомунікаційні інфраструктури без
значної зміни архітектури мережі** 88

Олег Іларіонов, Ганна Красовська, Григорій Гнатієнко,

Катерина Красовська, Равшанбек Зулунув

**Математична модель оцінювання збалансованості профілів
освітніх програм закладів вищої освіти** 92

Ярослав Дорогий, Владислав Кравчук, Василь Цуркан

**Тензорна модель представлення критичної
інфраструктури** 97

<i>Геннадій Шибачев, Леонід Гальчинський</i>	
Застосування методів машинного навчання для виявлення роботи кейлогерів в операційній системі	101
<i>Владислав Личик, Леонід Гальчинський</i>	
Апробація моделей кіберстійкості для функції розподілу енергосистеми	103
<i>Артем Попов</i>	
Архітектура нульової довіри	106
<i>Вячеслав Сенченко, Андрій Бойченко</i>	
Підхід до оцінки ризиків виникнення каскадних ефектів критичної інфраструктури	110
<i>Олександр Додонов, Анатолій Кузьмичов</i>	
Оптимально розподілена мережева структура розміщення і використання критичних ресурсів	118
<i>Віталій Циганок, Руслан Беселовський, Володимир Мінас, Віктор Голота, Олександр Григоренко</i>	
Розробка сценаріїв проблемних ситуацій із використанням автоматизованих рішень на платформі трансферу знань	120
<i>Дмитро Ланде, Анатолій Качинський</i>	
Графова модель ментальної війни	123
<i>Павло Роїк, Андрій Оленко, Сергій Каденко, Олег Андрійчук</i>	
Універсальний метод визначення рівня узгодженості експертних попарних порівнянь, достатнього для їх агрегації	129
<i>Quan Trong The</i>	
Likelihood Estimation - Based Post - Filtering for Generalized Sidelobe Canceller Beamformer	132
<i>Rexhep Mustafovski, Aleksandar Risteski, Tomislav Shuminoski</i>	
Biothreat early assist and response command system (BEAR-CS): a state-of-the-art analysis and comparative study with existing bio surveillance and incident management systems	134
<i>Андрій Снарський, Іван Дідур</i>	
Структурна складність порогових явищ у складних мережах	135
<i>Nataliia Kuznietsova, Illia Kvashuk, Petro Bidyuk</i>	
Survival models as copulas for green risks prediction	140
<i>Dmytro Kucherov, Xiaopeng Fang, Zhiqiang Wang, Zhongli Xu, Xiaojian Fu</i>	
Identification of parameters of the mathematical model of a dynamic object of critical infrastructure	143

<i>Oleksandr Pliushch, Bohdan Zhurakovskiy, Maryan Kyryk</i> Computer Model of Tuned Transistor Oscillator to Study its Excitation Processes	147
<i>Леся Ладієва, Богдан Корнієнко, Олександра Пилипон</i> Система керування процесом паро-кисневої конверсії метану	150
<i>Сергій Любарський, Едуард Бовда, Юрій Самохвалов</i> Оцінка знань у закладах вищої освіти на основі адаптивної IRT-моделі	154
<i>Ілля Куцербов, Павел Ломовцев</i> Геометричне моделювання фантастичної зброї для онлайн комп'ютерної гри у жанрі шутеру	157
<i>Єгор Загорій, Павел Ломовцев</i> Аналіз та порівняльне дослідження алгоритмів текстурування для реалістичних ігрових моделей	159
<i>Віталій Циганок, Ярослав Хроленко, Ірина Доманецька, Ольга Циганок, Олена Трубіцина</i> Автоматизований підхід до формування журі студентських наукових конкурсів на основі наукометричних та альтметричних показників	161
<i>Катерина Хала, Анатолій Гладун, Олександр Волков</i> Планування місій БпЛА на основі онтологій та рольового розподілу завдань у мультиагентній системі	165
<i>Андрій Бойченко</i> Використання технологій ГШІ для аналізу каскадних ефектів критичної інфраструктури	169
<i>Сергій Каденко, Олег Андрійчук, Галина Гоменюк</i> Моделювання неповних експертних оцінок, що враховують ранжирування	174
<i>Григорій Гнатієнко</i> Метод метризації вагових коефіцієнтів критеріїв при лексикографічному упорядкуванні альтернатив	177
<i>Ірина Храмова</i> Використання каталогів даних в управлінні безпекою корпоративних інформаційних ресурсів	181
<i>Вікторія Полуциганова, Сергій Смирнов</i> Уточнення оцінок вагових коефіцієнтів для методу парних порівнянь	185

<i>Вячеслав Корольов, Олег Рибальченко, Олександр Ходзінський, Максим Огурцов, Олександр Ярушевський</i>	
Визначення взаємної позиції дронів за умов відсутності GPS-сигналу для підтримки прийняття рішень під час виконання поставленої задачі	188
<i>Олексій Веретьонкін</i>	
Класифікація основних задач веб-застосунків	192
<i>Олексій Гайворонський, Василь Горбачук, Максим Дунаєвський, Сеїт-Бекір Сулейманов</i>	
Мережевий вплив даних на корпоративні цілі	193
<i>Микола Кіктєв, Роман Поліщук, Сергій Шворов, Вячеслав Іващук, Павло Мазурчук</i>	
Розумний енергозберігаючий промисловий тепличний комплекс з використанням технології IoT на базі Cisco Packet Tracer	197

Національна академія наук України
Інститут проблем реєстрації інформації НАН України

ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ТА БЕЗПЕКА

**Матеріали XXIV Міжнародної
науково-практичної конференції**

Випуск 24

Підп. до друку 26.12.2024. Формат 60х84¹/16. Папір офс. Гарнітура Times.
Спосіб друку – ризографія. Ум. друк. арк. 10,5. Обл.-вид. арк. 24,36. Наклад 100
пр. Зам. № 15-200.

ТОВ "Інжиніринг"
ISBN 978-617-8180-00-3