

Українська академія друкарства
Міністерство освіти і науки України
Національний університет «Львівська політехніка»
Міністерство освіти і науки України

Хмельницький національний університет
Міністерство освіти і науки України

Кваліфікаційна наукова праця
на правах рукопису

АНДРІЙВ РОМАН РОСТИСЛАВОВИЧ

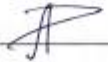
УДК 004.9:519.87:519.816

ДИСЕРТАЦІЯ
ІНФОРМАЦІЙНА ТЕХНОЛОГІЯ ОЦІНЮВАННЯ
ДОСТОВІРНОСТІ ТЕКСТОВИХ ПОВІДОМЛЕНЬ

05.13.06 – Інформаційні технології

Подається на здобуття наукового ступеня кандидата технічних наук

Дисертація містить результати власних досліджень. Подані до захисту наукові положення є власним напрацюванням. Всі використані ідеї, наукові результати, цитати супроводжуються належними посиланнями на їх авторів та джерела опублікування. Всі частини тексту дисертації, під час написання яких використовувалися технології штучного інтелекту, перевірені та відредаговані особисто.

_____  Андрійв Роман Ростиславович

Науковий керівник – Піх Ірина Всеволодівна, доктор технічних наук, професор.

*Ідентичність
примірників дисертації
засвідчую*

Хмельницький – 2025

Вчений секретар



Андрій НічЕТОРУК

АНОТАЦІЯ

Андрійв Р. Р. Інформаційна технологія оцінювання достовірності текстових повідомлень. Рукопис.

Дисертація на здобуття наукового ступеня кандидата технічних наук в галузі інформаційних технологій зі спеціальності 05.13.06 – Інформаційні технології. Хмельницький національний університет, Хмельницький, 2025.

У дисертаційній роботі розв’язано науково-прикладне завдання прогностного оцінювання достовірності текстових інформаційних повідомлень на основі аналізу факторів їх вірогідності шляхом розроблення інформаційної технології з використанням засобів теорії моделювання, дослідження операцій та нечітких множин.

Під час виконання роботи здійснено проєктування альтернативних підходів до реалізації процесу оцінювання достовірності текстових повідомлень та визначено його оптимальний варіант. Для формування інтегрального показника достовірності використано апарат нечіткої логіки, що дало змогу врахувати невизначеність і суб’єктивність вхідних даних. Результатом реалізації технології стало створення програмного застосунку, призначеного для прогностного оцінювання достовірності текстових інформаційних повідомлень (ДТП) до моменту поширення повідомлення у медійному просторі.

У *вступі* висвітлено актуальність теми дослідження з урахуванням сучасних тенденцій розвитку інформаційної галузі. Відзначено, що зростання соціальних мереж, онлайн-медіа та автоматизованих систем генерації контенту суттєво ускладнює перевірку достовірності повідомлень, що зумовлює потребу у впровадженні ефективних механізмів аналізу й оцінювання їхньої правдивості. Орієнтування інформаційних технологій на процеси оцінювання достовірності інформаційних повідомлень є актуальним завданням сучасного інформаційного суспільства.

Підкреслено зв’язок роботи з науковими дослідженнями кафедри комп’ютерних технологій у видавничо-поліграфічних процесах Інституту

поліграфії та медійних технологій, що орієнтуються на покращення якості процесів оцінювання достовірності інформаційних повідомлень. Визначено мету та основні завдання дослідження. Обґрунтовано наукову новизну та практичне застосування отриманих в процесі роботи результатів. Виокремлено особистий внесок здобувача в дисертацію та спільні публікації, що стосуються дослідження; описано науково-практичні конференції, на яких апробовано результати дисертації; надано підсумок стосовно структури та обсягу дисертації.

Перший розділ присвячено аналізу сучасного стану досліджуваної проблеми оцінювання достовірності текстових інформаційних повідомлень. Відзначено, що вивчення і дослідження вказаної тематики вимагає комплексного підходу для розв'язання складних проблем, пов'язаних з достовірністю інформації в сучасному інформаційному середовищі. Описано можливі напрями та варіанти реалізації результатів досліджень. Здійснено огляд систем та методів розпізнавання ступеня достовірності текстових повідомлень.

Наведено перелік теоретичних і прикладних ресурсів факторного аналізу, використаних для розроблення моделей і засобів процесу дослідження. До них віднесено: семантичні мережі; метод аналізу ієрархій та метод ранжування; метод попарних порівнянь; метод лінійного згортання критеріїв; модель нечіткого логічного виведення; функції належності лінгвістичних змінних, нечіткі матриці знань та нечіткі логічні рівняння. Розроблено багаторівневу модель виконання дисертаційного дослідження, орієнтовану на аналіз впливових факторів та оцінювання достовірності текстових інформаційних повідомлень.

Сформульовано основні завдання дисертаційної роботи.

У другому розділі виконано аналіз та опис факторів, що стосуються визначення показника достовірності текстових інформаційних повідомлень. Розроблено графічно орієнтовану модель у вигляді семантичної мережі, що уможливило відтворення формалізованого відображення структурних та лінгвістичних відношень між факторами. Виконано опис семантичної мережі з використанням елементів логіки предикатів.

Розроблено оптимізовану багаторівневу модель чинників достовірності текстових інформаційних повідомлень за методом математичного моделювання ієрархій на основі виконаного розрахунку вагових пріоритетів вказаних факторів та методу ранжування, що уможливило застосування методів аналізу ієрархій і теорії нечітких множин для подальшого дослідження. На підставі методу лінійного згортання критеріїв запроектовано альтернативні варіанти визначення рівня довіри до інформації на основі ефективності факторів множини Парето в альтернативах. Розраховано та визначено серед альтернативних оптимальний варіант процесу формування показника достовірності текстових даних за критерієм максимального значення цільового функціоналу.

Третій розділ присвячений теоретичному формуванню показника оцінювання прогностного рівня достовірності текстових інформаційних повідомлень. Охарактеризовано ключові поняття теорії нечітких множин, зокрема: лінгвістичні змінні (ЛЗ), функції належності (ФН), терм-множини значень, лінгвістичні терми факторів, бази нечітких знань, матриці знань, нечіткі логічні рівняння. Структуровано процес формування показника достовірності повідомлень, створено класифікаційні групи ЛЗ – організаційну, авторську та користувацьку.

Сформовано інформаційну базу даних, яка встановлює зв'язок між змінними, їхньою лінгвістичною природою, діапазонами значень універсальної терм-множини та визначеними лінгвістичними термами, заданими нечіткою тривимірною шкалою. Розроблено метод логічного виведення, багаторівнева модель яка відображає алгоритм формування показника достовірності текстових повідомлень з урахуванням класифікаційних груп лінгвістичних змінних.

Розроблено матриці попарних порівнянь для лінгвістичних термів кожної лінгвістичної змінної у п'яти точках поділу універсальної множини. Запроектовано нечіткі матриці знань, у яких відображено формалізовану реалізацію нечіткого логічного висновку, сформовано нечіткі логічні рівняння. Реалізовано процес формування показника прогностного рівня достовірності

текстових інформаційних повідомлень з використанням функцій належності, що базуються на попередньо встановлених значеннях лінгвістичних змінних.

У четвертому розділі проведено дефазифікацію нечіткої множини та експериментально розраховано значення показника достовірності текстових інформаційних повідомлень. Завершено процес розроблення структурно-функціональної моделі інформаційної технології прогностного оцінювання достовірності текстових повідомлень. Виконано імітаційне моделювання процесу оцінювання ДТІП із використанням методів багатокритеріальної оптимізації, що враховує множину впливових факторів та формалізує процедуру прийняття рішень щодо рівня достовірності повідомлень. Здійснено функціональне моделювання процесу створення інформаційної технології, реалізоване з використанням методології IDEF0, що дозволило відобразити логіку побудови технології у вигляді структурованих взаємозалежних процесів, уточнити рольових учасників, інформаційні потоки, вхідні й вихідні параметри, а також ресурси, необхідні для реалізації кожного з етапів.

Висновки узагальнено відображають основні результати та аналітичні підсумки дисертаційного дослідження.

У дисертації розв'язано науково-прикладне завдання прогностного оцінювання достовірності текстових інформаційних повідомлень на основі аналізу факторів їх вірогідності шляхом розроблення інформаційної технології з використанням засобів теорії моделювання, дослідження операцій та нечітких множин.

На основі виконаних досліджень отримано такі нові наукові результати:

- вперше розроблено моделі визначення оптимального варіанту процесу формування інтегрального показника оцінювання достовірності текстових інформаційних повідомлень, що забезпечує встановлення їх вірогідності;
- вперше розроблено метод логічного виведення з урахуванням класифікаційних груп лінгвістичних змінних – організаційної, технічної та користувацької – який відображає поетапну реалізацію процедур формування показника достовірності текстових інформаційних повідомлень.

- удосконалено метод формування показника достовірності текстових інформаційних повідомлень та встановлення їх вірогідності, у якому реалізовано фазифікацію як процес переходу від числових даних до нечітких оцінок на основі лінгвістичних змінних, а також дефазифікацію як зворотну процедуру перетворення нечітких результатів у конкретні числові значення.

- розроблено інформаційну технологію прогностного оцінювання рівня достовірності текстових повідомлень у частині використання засобів теорії моделювання, дослідження операцій та нечіткої логіки, що забезпечило визначення вірогідності інформаційного контенту.

Результати дисертаційної роботи використано в лекційних курсах та практичних заняттях Інституту поліграфії та медійних технологій Національного університету «Львівська політехніка», Львівського державного університету безпеки життєдіяльності, у видавничій діяльності.

Ключові слова: достовірність повідомлень, фактор, модель, матриця, нечітка логіка, лінгвістична змінна, терм-множина, функція належності, нечітка база знань, фазифікація, дефазифікація, інформаційна технологія, імітаційне моделювання.

ABSTRACT

Andriiv R. R. Information Technology for Assessing the Reliability of Text Messages. Manuscript.

Dissertation for the degree of Candidate of Technical Sciences in the field of Information Technologies, specialty 05.13.06 – Information Technologies. Khmelnytskyi National University, Khmelnytskyi, 2025.

In the dissertation, a scientific and applied problem of predictive assessment of the reliability of textual informational messages has been solved based on the analysis of their probability factors through the development of an information technology employing the tools of modeling theory, operations research, and fuzzy sets.

In the course of the research, alternative approaches to implementing the reliability assessment process were designed, and the optimal one was identified. To construct an integral indicator of reliability, fuzzy logic tools were employed, enabling the inclusion of uncertainty and subjectivity in input data. The result of implementing the technology was the development of a software application designed for predictive assessment of the reliability of textual informational messages (TIM) prior to their dissemination in the media space.

The introduction highlights the relevance of the research topic, taking into account current trends in the development of the information domain. It is noted that the growth of social networks, online media, and automated content generation systems significantly complicates the verification of message reliability, thus necessitating the implementation of effective mechanisms for assessing their truthfulness. The orientation of information technologies toward evaluating the reliability of information messages is identified as a pressing issue of the modern information society.

The work emphasizes its connection with the scientific research conducted by the Department of Computer Technologies in Publishing and Printing Processes at the Institute of Publishing and Media Technologies, which focuses on improving the quality of information credibility assessment processes. The objectives and key tasks

of the research are defined. The scientific novelty and practical applicability of the obtained results are substantiated. The author's individual contribution to the dissertation and relevant co-authored publications are outlined, as well as scientific and practical conferences where the findings were presented. A summary of the structure and scope of the dissertation is also provided.

Chapter One is devoted to the analysis of the current state of the research problem concerning the evaluation of the reliability of textual information messages. It is emphasized that addressing this issue requires a comprehensive approach to solving complex problems associated with the authenticity of information in today's digital environment. Possible directions and implementation strategies for the research findings are discussed. A review of existing systems and methods for identifying the degree of news credibility is presented.

The chapter also provides a set of theoretical and applied resources used for factor analysis in the design of models and tools within the study. These include semantic networks; the Analytic Hierarchy Process (AHP) and ranking methods; the method of pairwise comparisons; the linear convolution model of criteria; a fuzzy inference model; membership functions of linguistic variables; fuzzy knowledge matrices; and fuzzy logical equations. A multi-level model of the dissertation research process has been developed, focused on analyzing influential factors and assessing the reliability of textual information messages.

The core research tasks of the dissertation are formulated.

Chapter Two provides an in-depth analysis of the factors involved in the a priori determination of the reliability level of textual information messages. A graph-based semantic network model was developed to represent formalized relationships—both structural and linguistic—among these factors. This model was described using elements of predicate logic, enabling a structured semantic interpretation of interdependencies.

An optimized multi-level model of reliability factors was constructed using hierarchical mathematical modeling. The approach combined the calculation of factor

weight priorities with a ranking procedure, which allowed the application of the Analytic Hierarchy Process and fuzzy set theory in the subsequent stages of the study. Based on the method of linear convolution of criteria, alternative approaches to estimating the level of trust in information were proposed, relying on the efficiency of Pareto-optimal factor combinations. Among the generated alternatives, the optimal configuration for calculating the textual data reliability indicator was determined using the criterion of maximizing the target functional value.

Chapter Three is devoted to the theoretical construction of a predictive (a priori) indicator for evaluating the reliability of textual information messages. The chapter outlines fundamental concepts from fuzzy set theory, including linguistic variables (LV), membership functions (MF), value term-sets, linguistic descriptors of influencing factors, fuzzy knowledge bases, fuzzy knowledge matrices, and fuzzy logical equations. The process of indicator formation is structurally defined, with linguistic variables grouped into three classes: organizational, author-related, and user-related.

A structured database was developed to map the relationships between variables, their linguistic characteristics, corresponding value ranges of the universal term-set, and designated linguistic terms within a fuzzy three-dimensional scale. A multi-level fuzzy inference model was constructed to guide the process of generating the reliability indicator, incorporating the defined classes of linguistic variables.

Pairwise comparison matrices were developed for the linguistic terms of each variable across five partition points of the universal set. Fuzzy knowledge matrices were designed to formally implement the fuzzy inference mechanism, and fuzzy logical equations were constructed accordingly. The predictive reliability indicator was ultimately formed using membership functions derived from the preassigned values of the linguistic variables.

Chapter Four focuses on the defuzzification of the fuzzy set and the experimental calculation of the reliability indicator for textual information messages. The development of the structural and functional model of the information technology

for evaluating message reliability was finalized. Simulation modeling of the RTIM assessment process was carried out using multi-criteria optimization methods that account for a set of influencing factors and formalize the decision-making procedure regarding the degree of message reliability. Functional modeling of the information technology development process was conducted using the IDEF0 methodology, which enabled the representation of the system's logic as structured interrelated processes, clarification of role participants, information flows, input and output parameters, and the resources required at each implementation stage.

The conclusions summarize the main findings and analytical results of the dissertation study.

The dissertation addresses the scientific and applied task of predictive assessment of the credibility of textual information messages based on the analysis of news reliability factors by developing an information technology that utilizes modeling theory, operations research, and fuzzy sets.

Based on the conducted research, the following novel scientific results were obtained:

- for the first time, models have been developed to determine the optimal variant of the process of forming an integral indicator for assessing the reliability of textual informational messages, which ensures the establishment of their credibility.

- for the first time, a method of logical inference has been developed, taking into account the classification groups of linguistic variables – organizational, technical, and user-related – which reflects the step-by-step implementation of procedures for forming the reliability indicator of textual informational messages.

- the method for forming the reliability indicator of textual informational messages and establishing their credibility has been improved by implementing fuzzification as a process of transforming numerical data into fuzzy estimates based on linguistic variables, as well as defuzzification as the reverse procedure of converting fuzzy results into specific numerical values;

- an information technology for predictive evaluation of the reliability level of textual messages has been developed, involving the application of modeling theory, operations research, and fuzzy logic, which has enabled the determination of the credibility of informational content.

The results of the dissertation have been incorporated into lecture courses and practical classes at the Institute of Printing and Media Technologies of Lviv Polytechnic National University, the Lviv State University of Life Safety, as well as in publishing activities.

Keywords: message reliability, factor, model, matrix, fuzzy logic, linguistic variable, term set, membership function, fuzzy knowledge base, fuzzification, defuzzification, information technology, simulation modeling.

СПИСОК ОСНОВНИХ ПУБЛІКАЦІЙ ЗДОБУВАЧА

Статті у наукових виданнях, включених до Переліку наукових фахових видань України:

1. Піх І. В., Сеньківський В. М., Андріїв Р. Р. Проектування та розрахунок альтернативних варіантів реалізації технологічних процесів. *Технологія і техніка друкарства*. 2015 № 2 (48). С. 55-62. URL: [https://doi.org/10.20535/2077-7264.2\(48\).2015.48030](https://doi.org/10.20535/2077-7264.2(48).2015.48030).
2. Піх І. В., Сеньківський В. М., Андріїв Р. Р. Оптимізована модель чинників достовірності текстових даних. *Науковий вісник Національного лісотехнічного університету України*. 2024. Том 34 № 4. С. 78-85. URL: <https://doi.org/10.36930/40340410>. (Index Copernicus)
3. Піх І. В., Білик О. З., Сеньківський Н. Ю., Андріїв Р. Р., Браташ С. П. Багатофакторний вибір альтернативних варіантів процесу тестування ПЗ. *Вісник Вінницького політехнічного інституту*. 2024. № 3. С. 78-85. URL: <https://doi.org/10.31649/1997-9266-2024-174-3-78-85>. (Index Copernicus)
4. Гарасимчук О., Наконечний Ю., Луковський Т., Андріїв Р., Наконечний Т. Захищене зберігання даних із використанням технології блокчейн ETHEREUM. *Ukrainian Information Security Research Journal*. 2024. Том 26 №

1. С. 213-220. URL: <https://doi.org/10.18372/2410-7840.26.18845>.

5. Andriiv R., Pikh I. Method for ranking the reliability factors of text messages. *Computer Systems and Information Technologies*. 2025 (1), 178–184. URL: <https://doi.org/10.31891/csit-2025-1-21>.

6. Andriiv, R., Senkivskyy, V. Information technology for predicting the reliability level of text messages. *Computer Systems and Information Technologies*. 2025. (2), 65–71. URL: <https://doi.org/10.31891/csit-2025-2-7>.

Монографії (розділи у колективних монографіях):

7. Andriiv Roman. Modern methods of ensuring information protection in cybersecurity systems using artificial intelligence and Blockchain technology: collective monograph / O. Harasymchuk and others. Kharkiv: TECHNOLOGY CENTER PC, 2025. 132 p. (ISBN-13 (15) 978-617-8360-12-2) URL: <http://doi.org/10.15587/978-617-8360-12-2>

Публікації, які засвідчують апробацію матеріалів дисертації:

8. Піх І., Андріїв Р. Ідентифікація факторів впливу на якість проектування електронних видань. *Науково-технічна конференція професорсько-викладацького складу, наукових працівників і аспірантів: тези доповідей*, м. Львів, 16 лютого – 20 лютого 2015 р. / УАД Львів, 2015. С.118.

9. Піх І., Андріїв Р. Об'єктно-орієнтована модель процесу проектування електронних видань. *Науково-технічна конференція професорсько-викладацького складу, наукових працівників і аспірантів: тези доповідей*, м. Львів, 16 лютого – 19 лютого 2016 р. / УАД, Львів, 2016. С. 145.

10. Беспалько О., Ткачук Р., Андріїв Р. Дослідження методів захисту інформації веб-сайтів на основі моделей розподілення доступу та моніторингу ідентифікаторів користувача. Зб. тез доп. *VI Всеукр. наук.-практ. конф. молодих учених, студентів і курсантів «Інформаційна безпека та інформаційні технології»*, м. Львів, 30 листопада 2023 р. / ЛДУБЖД Львів, 2023. С. 16-20.

11. Pikh I., Senkivsky V., Kudriashova A., Tupychak L., Andriiv R. Formation of an indicator of the information message reliability level by fuzzy logic toolkit. *CEUR Workshop Proceedings*, 2024, Vol. 3899, P. 99-112. URL: <https://ceur-ws.org/Vol-3899/paper10.pdf> (індексована в наукометричній базі Scopus)

12. Senkivsky V., Pikh I., Kudriashova A., Andriiv R., Senkivskyi N.. Evaluation of methods for intellectual analysis of manipulative content in the information space. *CEUR Workshop Proceedings*, 2025, Vol. 3963, P. 139-155. URL: <https://ceur-ws.org/Vol-3963/paper12.pdf> (індексована в наукометричній базі Scopus)

Публікації, які додатково відображають наукові результати дисертації:

13. Сеньківський В. М., Піх І. В., Андріїв Р. Р. Синтез моделі факторів композиційного оформлення іміджевої презентації. *Поліграфія і видавнича справа*. 2011. № 4 (56), С. 117-124. (Index Copernicus)

14. Андріїв Р. Р., Піх І. В. Вибір альтернативної системи для формування навчальних електронних ресурсів. *Поліграфія і видавнича справа*. 2015. № 1 (69), С. 69-76. (Index Copernicus)

15. Піх І. В., Сеньківська Н. Є., Білик О. З., Сеньківський Н. Ю., Андріїв Т. Р., Андріїв Р. Р. Модель факторів якості тестування програмного забезпечення. *Комп'ютерні технології друкарства*. 2024 № 1(51). С. 90-108. URL: <http://doi.org/10.32403/2411-9210-2024-1-51-90-108>.

16. Піх І. В., Андріїв Р. Р., Коцік О. Ю. (2024). Інноваційні підходи до розробки та впровадження стартапів. *Науково-технічна конференція професорсько-викладацького складу, наукових працівників і аспірантів: тези доповідей*. м. Львів, 5 лютого – 9 лютого 2024 р./ УАД Львів, 2024. С.186.

ЗМІСТ

ВСТУП.....	16
РОЗДІЛ 1. АНАЛІЗ МЕТОДІВ, ЗАСОБІВ І ТЕХНОЛОГІЙ ОЦІНЮВАННЯ ДОСТОВІРНОСТІ ТЕКСТОВИХ ІНФОРМАЦІЙНИХ ПОВІДОМЛЕНЬ	23
1.1. Аналіз проблематики оцінювання достовірності	
текстових інформаційних повідомлень.....	23
1.2. Сучасний стан досліджень тематики.....	
оцінювання достовірності текстових повідомлень	25
1.3. Огляд систем визначення рівня достовірності текстових повідомлень та правові аспекти використання	30
1.4. Засоби факторного аналізу для дослідження проблематики	
оцінювання рівня достовірності текстових повідомлень	37
1.6. Завдання дисертаційної роботи	45
Висновки до розділу 1	46
РОЗДІЛ 2. МОДЕЛІ ФАКТОРІВ ВПЛИВУ НА ДОСТОВІРНІСТЬ ТЕКСТОВИХ ІНФОРМАЦІЙНИХ ПОВІДОМЛЕНЬ	48
2.1. Фактори процесу формування рівня	
достовірності текстових повідомлень	48
2.2. Формалізоване відтворення зв'язків між факторами	
засобами семантичних мереж та мови предикатів	52
2.3. Модель ієрархічного впливу факторів на формування	
показника достовірності текстових інформаційних повідомлень.....	56
2.4. Модель факторів достовірності текстових повідомлень на основі методу ранжування та методу визначення вагомості предикатів	61
2.5. Оптимізація моделі факторів формування показника достовірності текстових інформаційних повідомлень.....	68
2.6. Альтернативні та оптимальний варіанти процесу	
оцінювання рівня достовірності текстових повідомлень	75
Висновки до розділу 2.....	85
РОЗДІЛ 3. МЕТОДИ ФОРМУВАННЯ ПОКАЗНИКА ОЦІНЮВАННЯ ДОСТОВІРНОСТІ ТЕКСТОВИХ ІНФОРМАЦІЙНИХ ПОВІДОМЛЕНЬ	87
3.1. Основні поняття і засоби теорії нечіткої логіки	87
3.2. Проектування універсальної терм-множини значень.....	93
3.3. Метод нечіткого логічного виведення як основа формування показника	

рівня достовірності повідомлень	97
3.4. Функції належності лінгвістичних змінних як чинників формування показника оцінювання рівня достовірності текстових повідомлень	99
3.5. Проектування нечітких матриць знань та нечітких логічних рівнянь.....	117
Висновки до розділу 3	123
РОЗДІЛ 4. ІНФОРМАЦІЙНА ТЕХНОЛОГІЯ ОЦІНЮВАННЯ ДОСТОВІРНОСТІ ТЕКСТОВИХ ПОВІДОМЛЕНЬ	125
4.1. Експериментальний розрахунок показника	126
4.2. Оцінювання рівнів достовірності текстових інформаційних повідомлень методом багатокритеріальної оптимізації	133
4.3. Функціональне моделювання процесу розроблення інформаційної технології оцінювання достовірності текстових повідомлень	139
4.4. Структурно-функціональна модель інформаційної технології оцінювання достовірності текстових повідомлень	144
4.5. Порівняння методів оцінювання достовірності текстових повідомлень	149
Висновки до розділу 4	151
ВИСНОВКИ	154
СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ	157
ДОДАТОК А. ПЕРЕЛІК ПУБЛІКАЦІЙ ЗДОБУВАЧА	170
ДОДАТОК Б. ДОВІДКИ ПРО ВИКОРИСТАННЯ В НАВЧАЛЬНОМУ ПРОЦЕСІ РЕЗУЛЬТАТІВ ДИСЕРТАЦІЙНОЇ РОБОТИ.....	173
ДОДАТОК В. АКТ ВИПРОБУВАНЬ ПРОГРАМНОГО ЗАСТОСУНКУ ВИЗНАЧЕННЯ ПРОГНОЗНОГО РІВНЯ ОЦІНЮВАННЯ ДТІП.....	175
ДОДАТОК Г ПРОГРАМНИЙ КОД	176

ВСТУП

Актуальність теми. У контексті сучасного інформаційного простору, що характеризується високою динамікою цифровізації, масовим використанням мережевих ресурсів та інтенсивним поширенням контенту через соціальні платформи, проблема оцінювання достовірності текстових інформаційних повідомлень набуває особливої актуальності. З огляду на значне зростання обсягів інформаційних потоків та складність верифікації їх джерел, постає об'єктивна потреба у створенні дієвих інструментів для виявлення рівня правдивості, надійності та змістовної обґрунтованості повідомлень, що циркулюють у відкритих комунікаційних середовищах.

Останніми роками тематика достовірності інформації привертає увагу широкого кола дослідників, які застосовують міждисциплінарний підхід, поєднуючи здобутки комп'ютерних наук, психології, соціології та медіа освіти. У межах сучасних наукових підходів вивчаються різні аспекти проблеми, зокрема обробка природної мови та аналіз фейкових новин, поширення дезінформації в соціальних мережах, особливості наукової та політичної комунікації, а також інформаційні війни. Дослідження у цих напрямках активно ведуться такими зарубіжними науковцями, як M. Castells [3], S. Mishra [4], D. Mikhail [7], S. Rastogi [8], B. Rath [9], M. Basol [12], G. Rampersad [13], S. Jones-Jang [14], D. Caled [16]. Серед українських вчених суттєвий внесок у розвиток зазначеної тематики здійснили Б. Дурняк [46], Т. Говорущенко [62], О. Бармак [37], М. Коробчинський [23], Р. Ткачук [48], І. Піх [50], В. Сабат [38], А. Кудряшова [44], В. Сеньківський [40] та ін.

Розвитком наукового доробку у даній галузі є задекларований та реалізований в дисертаційній роботі перспективний напрям досліджень, який полягає у застосуванні факторного аналізу для розв'язання проблеми оцінювання достовірності інформаційних повідомлень. Зокрема, враховано множину чинників, що впливають на достовірність даних, запропоновано новий підхід до формування та розрахунку показника прогнозного рівня вірогідності

повідомлень із використанням методів теорії моделювання та нечіткої логіки.

Таким чином, прогнозне оцінювання достовірності текстових інформаційних повідомлень на основі факторного аналізу їх вірогідності із застосуванням інформаційної технології, розробленої на засадах теорії моделювання, дослідження операцій та апарату нечітких множин, є актуальним науково-практичним завданням.

Сформульоване завдання відповідає паспорту спеціальності 05.13.06 (зокрема, п. 2 і п. 9).

Зв'язок роботи з науковими програмами, планами, темами. Тема дисертації пов'язана з науковими та прикладними дослідженнями кафедри комп'ютерних технологій у видавничо-поліграфічних процесах, відповідає науковому напрямку «Дослідження, розроблення і впровадження інформаційних технологій та систем для забезпечення достовірності інформації» Національного університету «Львівська політехніка».

Дисертація виконана в межах науково-дослідної роботи Української академії друкарства «Проектування реклами в мережі Інтернет на основі її семантичного аналізу» (договір № 601-2023 від 01.01.2023 року), в якій автор був виконавцем.

Мета і задачі дослідження. Метою дисертаційної роботи є прогнозне оцінювання достовірності текстових інформаційних повідомлень на основі аналізу факторів їх вірогідності шляхом розроблення інформаційної технології з використанням засобів теорії моделювання, дослідження операцій та нечітких множин.

Для вирішення вказаної проблематики необхідне розв'язання таких завдань:

- виконати аналіз методів, засобів та технологій оцінювання рівня достовірності інформаційних повідомлень;
- розробити моделі визначення оптимального варіанту процесу формування інтегрального показника оцінювання достовірності текстових

інформаційних повідомлень;

- розробити метод логічного виведення та сформуванати нечіткі матриці знань та нечіткі логічні рівняння для лінгвістичних рівнів моделі;

- удосконалити метод формування показника достовірності текстових повідомлень та визначення їх вірогідності, у якому реалізовано фазифікацію як процес переходу від числових даних до нечітких оцінок на основі лінгвістичних змінних;

- виконати дефазифікацію нечітких множини, реалізовану на основі нечіткого логічного виведення, що забезпечить отримання для заданих варіантів вихідних даних показника прогнозованого рівня достовірності текстових повідомлень;

- розробити інформаційну технологію прогнозного оцінювання достовірності текстових інформаційних повідомлень з використанням засобів теорії моделювання, дослідження операцій та нечітких множин.

Об'єкт дослідження – процеси аналізу та оцінювання достовірності текстових інформаційних повідомлень.

Предмет дослідження – моделі, методи та засоби інформаційної технології оцінювання рівня достовірності текстових інформаційних повідомлень.

Методи досліджень. У дисертаційному дослідженні реалізовано міждисциплінарний підхід із залученням низки формальних, аналітичних і прикладних методів. Зокрема, системний і матричний аналіз, положення теорії мереж та елементи логіки предикатів використано для формалізації взаємозв'язків між факторами, що впливають на рівень достовірності інформаційних повідомлень. Застосування методів моделювання ієрархічних структур та процедур ранжування дозволило сконструювати узагальнені моделі пріоритетності впливу зазначених факторів на цільові процеси оцінювання. Теоретичні засади дослідження операцій були залучені до задач оптимізації моделей, що передбачали обґрунтований вибір між множиною альтернатив та

ідентифікацію оптимального сценарію реалізації відповідних аналітичних дій. Ключову роль у побудові механізму оцінювання відіграли засоби теорії нечітких множин і нечіткої логіки, які стали основою для розроблення ієрархічної структури нечіткого логічного виведення, формування бази знань, побудови системи нечітких рівнянь, а також розрахунку функцій належності лінгвістичних змінних. Впровадження інформаційної технології формування прогнозного показника достовірності стало можливим завдяки використанню програмно-інженерних засобів, на основі яких реалізовано прикладне рішення у вигляді програмного модуля автоматизованого оцінювання вірогідності текстових повідомлень.

Наукова новизна отриманих результатів. На підставі теоретичних і практичних досліджень, виконаних у дисертаційній роботі, отримано такі нові наукові результати:

- вперше розроблено моделі визначення оптимального варіанту процесу формування інтегрального показника оцінювання достовірності текстових інформаційних повідомлень, що забезпечує встановлення їх вірогідності;

- вперше розроблено метод логічного виведення з урахуванням класифікаційних груп лінгвістичних змінних – організаційної, технічної та користувацької – який відображає поетапну реалізацію процедур формування показника достовірності текстових інформаційних повідомлень.

- удосконалено метод формування показника достовірності текстових інформаційних повідомлень та встановлення їх вірогідності, у якому реалізовано фазифікацію як процес переходу від числових даних до нечітких оцінок на основі лінгвістичних змінних, а також дефазифікацію як зворотну процедуру перетворення нечітких результатів у конкретні числові значення.

- розроблено інформаційну технологію прогнозного оцінювання рівня достовірності текстових повідомлень у частині використання засобів теорії моделювання, дослідження операцій та нечіткої логіки, що забезпечило визначення вірогідності інформаційного контенту.

Практичне значення отриманих результатів. У дисертаційній роботі розроблено методи, моделі та інформаційну технологію оцінювання достовірності текстових повідомлень.

Практично вагомими вважаються такі результати:

- описано основні етапи процесу класифікації текстових даних;
- розроблено оптимізовані моделі пріоритетного впливу виокремлених факторів на процес формування показника рівня достовірності текстових інформаційних повідомлень;
- розроблено метод логічного виведення, що відображає логіку формування показника достовірності текстових інформаційних повідомлень відповідно до виділених груп лінгвістичних змінних;
- удосконалено метод формування показника прогнозованого рівня достовірності текстових повідомлень;
- розроблено інформаційну технологію оцінювання рівня достовірності текстових інформаційних повідомлень, що визначає ступінь їх правдивості.

Результати дисертаційної роботи апробовано в умовах друкарні ТОВ «Видавничий Дім «Високий замок»» (м. Львів), використано у навчальному процесі Львівського державного університету безпеки життєдіяльності та Інституту поліграфії та медійних технологій НУ «Львівська політехніка». Дані про впровадження підтверджено відповідними документами.

Особистий внесок здобувача. Усі наукові результати дисертаційної роботи, що виносяться на захист, отримані автором особисто. У спільних публікаціях внесок здобувача виражається таким чином: [1] – експериментальний розрахунок варіантів виконання технологічних процесів; [2] – модель чинників достовірності текстових даних; [3] – вибір альтернативних варіантів процесу тестування ПЗ; [4] – аналіз процесу захищеного зберігання даних; [5] – метод ранжування факторів достовірності текстових повідомлень; [6] – інформаційна технологія прогнозування рівня достовірності текстових повідомлень; [7] – підбір публікацій стосовно сучасних методів забезпечення

захисту інформації; [8] – ідентифікація факторів впливу на якість проектування електронних видань; [9] – модель процесу проектування електронних видань; [10] – методи захисту інформації веб-сайтів; [11] – формування засобами нечіткої логіки показника рівня достовірності інформаційного повідомлення; [12] – оцінювання методів інтелектуального аналізу маніпулятивного контенту; [13] – модель факторів композиційного оформлення презентації; [14] – вибір альтернативної системи для формування навчальних електронних ресурсів; [15] – модель факторів якості тестування ПЗ; [16] – інноваційні підходи до розроблення та впровадження стартапів.

Внесок інших співавторів у спільних публікаціях: у статті [1] І. Піх формувала концепцію дослідження, В. Сеньківський сприяв підготовці плану дослідження та попереднього варіанту рукопису; у статті [2] І. Піх разом із В. Сеньківським виконували концептуалізацію дослідження; у статті [3] І. Піх обґрунтувала застосування загальних варіантів тестування ПЗ до процесу розроблення і тестування програми визначення рівнів достовірності повідомлень, О. Білик зумовив підготовку рукопису, Н. Сеньківський та С. Браташ допомагали підготувати чернетку рукопису; у статті [4] Ю. Наконечний обумовив концепцію дослідження, О. Гарасимчук складав план дослідження, Т. Луковський сприяв підготовці рукопису, Т. Наконечний забезпечив дані стосовно технології блокчейн; у статті [5] І. Піх формувала концепцію дослідження; у статті [6] В. Сеньківський виконував концептуалізацію і план дослідження; у монографії [7] О. Гарасимчук і Т. Наконечний виконували обґрунтування дослідження та формування структури монографії; у публікаціях [8,9] І. Піх формувала концепцію дослідження; у публікації [10] О. Беспалько спільно з Р. Ткачуком виконували концептуалізацію дослідження; у статті [11] І.Піх формувала концепцію дослідження, А. Кудряшова складала план дослідження, Л. Тупичак виконувала візуалізацію результатів дослідження; у статті [12] В. Сеньківський формував концепцію і план дослідження, І. Піх спільно з А. Кудряшовою сприяли підготовці рукопису, Н. Сеньківський

виконував візуалізацію дослідження; у статті [13] В. Сеньківський формував план дослідження, І. Піх сприяла концептуальним рішенням; у статті [14] І. Піх розробляла план дослідження; у статті [15] І. Піх виконувала концептуалізацію дослідження, О. Білик та Н. Сеньківська сприяли підготовці рукопису, Н. Сеньківський і Т. Андріїв відповідали за апробацію розробленої моделі; у публікації [16] І. Піх формувала тематику дослідження, О. Коцік виконував підготовку вхідних даних.

Апробація результатів дисертації. Результати дисертації доповідали та обговорювали на: звітних науково-технічних конференціях професорсько-викладацького складу, наукових працівників і аспірантів Української академії друкарства (Львів, 2015, 2016, 2024); VI-й Всеукр. наук.-практ. конф. молодих учених, студентів і курсантів «Інформаційна безпека та інформаційні технології» у Львівському державному університеті безпеки життєдіяльності (Львів, 2023); The 1st International Workshop on Advanced Applied Information Technologies, (Khmelnyskyi, 2025); The 6th International Workshop on Intelligent Information Technologies & Systems of Information Security, (Khmelnyskyi, 2025).

Публікації. За результатами досліджень опубліковано 16 наукових праць, до яких відносяться: 8 публікацій у наукових журналах, що входять до міжнародної наукометричної бази даних Index Copernicus, та є науковими фаховими виданнями України; одна монографія; 7 публікацій, які засвідчують апробацію матеріалів дисертації (серед яких дві статті у матеріалах конференцій, що індексуються в наукометричній базі Scopus).

Структура та обсяг дисертації. Дисертація складається з анотації, вступу, чотирьох розділів, висновків, списку використаних джерел із 96 найменувань і трьох додатків. Загальний обсяг рукопису становить 176 сторінок, в тому числі 156 сторінок основного тексту, 19 рисунків та 36 таблиць.

РОЗДІЛ 1. АНАЛІЗ МЕТОДІВ, ЗАСОБІВ І ТЕХНОЛОГІЙ ОЦІНЮВАННЯ ДОСТОВІРНОСТІ ТЕКСТОВИХ ІНФОРМАЦІЙНИХ ПОВІДОМЛЕНЬ

1.1. Аналіз проблематики оцінювання достовірності текстових інформаційних повідомлень

Проблематика оцінювання достовірності текстових інформаційних повідомлень є надзвичайно важливою та актуальною в умовах сучасного інформаційного суспільства. Стрімкий розвиток цифрових технологій, поширення Інтернету та соціальних медіа, а також зростання обсягів генерованої інформації зумовлюють необхідність розробки ефективних підходів до визначення правдивості, достовірності та якості отриманих даних.

Актуальність цього питання також зумовлена його безпосереднім впливом на інформаційну безпеку, рівень медіа грамотності та стабільність функціонування суспільства. Вирішення зазначених проблем потребує комплексного підходу, що поєднує методи аналізу текстових даних, штучного інтелекту та обчислювальної лінгвістики, що дає змогу ефективно оцінювати достовірність інформації в сучасному інформаційному просторі.

Боротьба з фейковими новинами є надзвичайно складним завданням, оскільки в демократичних суспільствах свобода слова є фундаментальним правом, що сприяє автономності та плюралізму засобів масової інформації. Одночасно спостерігається складність у розрізненні між проявом індивідуальної, часто нетрадиційної, точки зору та свідомим розповсюдженням хибної інформації, що значно ускладнює ідентифікацію достовірних повідомлень у загальному інформаційному потоці..

Оцінювання достовірності текстових повідомлень та новинних матеріалів може здійснюватися за кількома основними напрямками. Одним із ключових аспектів є визначення авторства тексту, аналіз репутації джерела, його історії та рівня впливовості. У цьому контексті важливо не лише встановити факт існування джерела, а й оцінити його наміри щодо поширення інформації, що

може свідчити про можливі маніпулятивні або дезінформаційні стратегії.

Перевірка фактів та подій, описаних у тексті, є важливим аспектом оцінювання достовірності інформації. Цей процес може включати залучення фактчекінгових методів, порівняння зі змістом незалежних джерел, аналіз контенту та обставин, за яких було створено повідомлення. Важливим вважається виявлення спроб маніпуляції аудиторією через поширення неправдивої або спотвореної інформації. Такі маніпуляції можуть мати різну мотивацію — політичну, соціальну, економічну або іншу. Аналіз змісту повідомлень з урахуванням їхнього стилю, консистентної та логічної узгодженості дозволяє оцінити загальне враження від отриманої інформації, що є важливим показником її правдивості.

Використання алгоритмів машинного навчання забезпечує можливість автоматичної оцінки достовірності текстів на основі великої кількості показників. Такі алгоритми можуть враховувати семантичні, синтаксичні та статистичні характеристики тексту, а також метаінформацію про джерело.

Дослідження також охоплюють питання впливу соціальних мереж на поширення інформації та механізми контролю її достовірності. Алгоритми, що використовуються платформами для модерації контенту, можуть як сприяти, так і ускладнювати боротьбу з дезінформацією.

Важливим аспектом є аналіз сприйняття інформації громадськістю та факторів, що впливають на рівень довіри до різних джерел. Розробка стратегій підвищення довіри до якісних медіа сприятиме формуванню критичного мислення та протидії поширенню фейкових новин.

Оцінювання достовірності текстових повідомлень може здійснюватися на різних рівнях узагальнення — від аналізу окремих випадків до виявлення глобальних тенденцій у поширенні дезінформації. Це також включає визначення відповідальності за перевірку фактів, розробку стандартів професійної етики та роль журналістів і медійних організацій у забезпеченні достовірності інформації.

Етичні питання, пов'язані з оцінкою достовірності текстових повідомлень,

охоплюють проблеми конфіденційності джерел, баланс між інформаційною відкритістю та захистом суспільства від дезінформації.

Фейкові новини можуть бути виявлені через багатовимірний аналіз їхньої узгодженості з іншими даними. Це включає перевірку технічної інформації щодо справжності джерела, аналіз соціального та/або юридичного бекграунду повідомлень, а також оцінку їхньої відповідності загальновизнаним фактам. Успішна перевірка фактів вимагає глибокого розуміння контенту та використання надійних джерел інформації. Великий обсяг фейкових новин, що поширюються через соціальні мережі, унеможлиблює їхню перевірку виключно вручну. Хоча ручна перевірка фактів може бути ефективною у деяких випадках, наприклад, під час аналізу узгодженості текстових повідомлень у різних контекстах, її швидкість є недостатньою для обробки значної кількості інформації, що поширюється у цифровому середовищі.

Сучасні системи автоматизованої верифікації інформації ґрунтуються на багаторівневому підході, який інтегрує технології штучного інтелекту, алгоритми обробки природної мови та інноваційні механізми на основі блокчейн-рішень з метою підвищення точності та ефективності виявлення недостовірного контенту. У разі фейкових новин, вбудованих у зображення та відео, такі системи аналізують широкий спектр характеристик, включаючи метадані, соціальні взаємодії, візуальні підказки, профіль джерела, а також інші дані, що супроводжують мультимедійний контент. Використання подібних методів сприяє підвищенню точності виявлення маніпулятивної або недостовірної інформації.

1.2. Сучасний стан досліджень тематики оцінювання достовірності текстових повідомлень

Ознайомлення з особливостями предметної області передбачає аналіз наукових праць, що відтворюють результати досліджень процесів забезпечення

достовірності текстових даних. Оскільки актуальність озвученої проблеми обумовила наявність значної кількості публікацій, наведено деякі з них, найбільш близькі до тематики статті.

Важливі аспекти ролі текстових повідомлень в різних сферах і напрямках людської діяльності наведено в роботах [1-4]. Автори публікацій підкреслюють, що текстові повідомлення та інформація є не просто інструментами, але й критичними елементами для розвитку та функціонування сучасного суспільства. Їх ефективне використання може привести до суттєвого покращення у багатьох сферах людської діяльності.

Одним із вагомих джерел, що має концептуальний зв'язок із тематикою дисертаційного дослідження, є публікація [5], у якій запропоновано методологічну модель виявлення фейкових новин на основі чітко структурованого багатофазного підходу. Автори розглядають чотири послідовні етапи: попередню обробку даних, редукцію ознак, їхню детальну екстракцію та подальшу класифікацію. Застосування регіональної розрахункової системи автоматизації (PPCA) дає змогу підвищити точність обробки вхідної інформації за рахунок оптимізації параметрів, отриманих на попередніх етапах. Ключовим компонентом етапу класифікації є використання архітектури LSTM (Long Short-Term Memory) — різновиду рекурентної нейронної мережі, ефективною для аналізу послідовних структур, зокрема текстових повідомлень. У результаті реалізації зазначеного підходу формується система, здатна з високим рівнем точності визначати ступінь достовірності інформаційних матеріалів.

У роботі [6] представлено оригінальний підхід до задачі класифікації недостовірної інформації, що суттєво відрізняється від традиційних моделей. Автори розглядають концепцію Elementary Discourse Unit (EDU), яка зосереджується на аналізі рівня деталізації між окремими словами та реченнями в тексті, що дозволяє глибше усвідомити структуру повідомлення. Запропонована модель класифікації фейкових новин базується на цій детальній лінгвістичній сегментації, що сприяє підвищенню точності розпізнавання

маніпулятивного контенту. Результати дослідження демонструють, що використання підходу EDU забезпечує значне вдосконалення існуючих систем верифікації інформації

У роботі [7] запропоновано інноваційний підхід до виявлення фейкових новин, який базується на моделюванні графів поширення інформації. Ця методика використовує аналіз профілів користувачів соціальних мереж, їхніх взаємозв'язків і характеру розповсюдження контенту, що значно підвищує надійність виявлення недостовірних повідомлень і розширює можливості застосування у різних інформаційних середовищах. На відміну від традиційних підходів, що зосереджені виключно на текстовій класифікації, описаний метод дозволяє враховувати соціальну динаміку поширення фейків. Дослідження [8] представляє комплексний фреймворк, здатний розпізнавати інформацію у різних соціальних мережах та класифікувати її за типами, включно з достовірними даними, дезінформацією і сатиричним контентом. У публікації детально описані базові принципи побудови фреймворку, алгоритми класифікації та процес ідентифікації контенту. Робота [9] розглядає застосування графічних нейронних мереж для ідентифікації вузлів, які потенційно можуть бути джерелами розповсюдження фейкової інформації, що дозволяє формувати прогнози щодо поведінки користувачів у мережі. Особливістю цього підходу є аналіз ефективності системи у присутності ботів, що активно поширюють неправдиві дані, та здатність виявляти таких розповсюджувачів ще до того, як їхній контент пройде перевірку на достовірність. Сучасні дослідження у цій галузі спрямовані на розробку алгоритмічних і аналітичних рішень для виявлення маніпулятивних ознак, аналізу джерел інформації та визначення ступеня її надійності [10,11]. Використання методів машинного навчання та штучного інтелекту дозволяє підвищити точність оцінки інформаційних потоків, мінімізуючи негативний вплив деструктивного контенту на суспільство [12]. Активний розвиток технологій штучного інтелекту, лінгвістичного аналізу та нейронних мереж забезпечує поступове вдосконалення засобів обробки як текстових, так і

мультимедійних даних [13,14]. У роботі [15] здійснюється критичний аналіз ефективності сучасних методів штучного інтелекту і машинного навчання щодо виявлення загроз і аномалій у великих обсягах інформації. Розглянуто різні підходи, їхні переваги та недоліки та перспективи розвитку.

Особливістю досліджуваної нами проблематики є висвітлення підходів у праці [16], які об'єднують обчислювальні рішення та пропонують стратегії боротьби з дезінформацією, зокрема, використання методик на основі машинного навчання для автоматичної класифікації дезінформації. Методи машинного навчання використовують статистичні підходи [17,18]. для ідентифікації закономірностей та аномальних відхилень у великих масивах даних, що дозволяє аналітикам безпеки своєчасно виявляти потенційні загрози, зокрема раніше невідомі атаки.

Оцінка ефективності подібних методів може ґрунтуватися на принципі Парето, який дозволяє обирати оптимальні рішення за критеріями точності, швидкодії та ресурсозатратності. Наприклад, моделі глибокого навчання [19], демонструють високу точність, однак потребують значних обчислювальних ресурсів, тоді як класичні методи машинного навчання менш затратні, проте можуть поступатися у здатності розпізнавати складні маніпулятивні конструкції. Обрання певної методики аналізу визначається специфікою інформаційного середовища, в якому проводиться дослідження, а також вимогами до рівня точності та оперативності обробки даних [20].

Наукові праці вітчизняних дослідників, присвячені цій проблематиці, системно висвітлюють і критично аналізують актуальні виклики та загрози, що виникають у сфері інформаційної безпеки України. Особлива увага приділена вивченню механізмів протидії інформаційно-психологічним впливам, які мають безпосередній вплив на національну стійкість та безпеку [21]. Звернено увагу на важливість вміння оцінювати достовірність текстових даних та загальної інформації про певні події, можливість її появи в певних часових і просторових межах, заданих наявним інформаційним описом, здатність визначення рівня

достовірності інформації як міри загрози для інформаційного простору [22]. Наведено перелік питань, на які треба звернути увагу при аналізі достовірності інформації та визначенні її правдивості [23]. Для оцінювання достовірності даних в інформаційному просторі рекомендовано звернути увагу на певних ознаках, до яких віднесено сумнівність викладених фактів, емоційність, тональність та сенсаційність контенту [24]. Пропоновано критерії моніторингу достовірності текстових даних в інформаційному просторі [25], суть яких зводиться до балансу інформації, відокремлення фактів від загальних міркувань, окреслення точності, об'єктивності та повноти інформації. Описаний у публікації [26] інформаційний підхід ґрунтується на виокремленні факторів впливу на показник достовірності текстових даних. Виконано формалізоване відтворення зв'язків між факторами за допомогою семантичної мережі, опрацювання якої обумовило ранжування факторів за умовними ваговими пріоритетами, що виступає початковим етапом більш повного дослідження, яке передбачає розроблення інформаційної технології прогностного оцінювання достовірності текстових повідомлень.

Проведений аналіз наукових джерел, навіть за умови охоплення лише обмеженого кола публікацій, підтверджує надзвичайну актуальність розроблення ефективних засобів протидії спотворенню інформаційних повідомлень – незалежно від їхньої форми: текстової, усної чи електронної. З огляду на зростання обсягів дезінформації в медійному просторі, особливої ваги набуває завдання своєчасного виявлення фейкових новин з метою зменшення їх деструктивного впливу на суспільну свідомість. Запропонований у дослідженні підхід до прогностного оцінювання достовірності текстових повідомлень ще до їх появи забезпечує дієвий інструмент для попередження поширення недостовірної інформації, сприяючи формуванню безпечнішого інформаційного середовища для користувачів.

1.3. Огляд систем визначення рівня достовірності текстових повідомлень та правові аспекти використання

У видавничій та медійній сфері, де інформаційна точність та якість контенту є ключовими критеріями, системи визначення рівня достовірності текстових повідомлень набувають особливого значення. З огляду на зростання обсягів цифрового контенту та загрозу дезінформації, видавництва дедалі частіше інтегрують інтелектуальні інструменти перевірки достовірності в редакційні процеси. Такі системи поєднують методи обробки природної мови (NLP), машинного навчання, семантичного аналізу та логічного виводу, що дозволяє оцінювати відповідність тексту перевіреним джерелам, виявляти ознаки маніпуляцій та визначати ступінь фактологічної відповідності [25]. Медійна сфера вирізняється високою динамікою оновлення інформації, великою кількістю неструктурованих текстів та потребою в оперативності. У таких умовах особливо актуальними стають системи, здатні не лише оцінювати достовірність фактів, а й враховувати контент повідомлення, його тональність, джерело походження, ситуативну актуальність та навіть потенційний вплив на суспільство [27-29].

Зокрема, у світовій практиці поширення набули системи типу ClaimBuster, Full Fact, AdVerif.ai, які орієнтовані на виявлення фактологічних похибок та неузгодженостей на рівні речень або абзаців. Їх алгоритмічна основа ґрунтується на поєднанні навчання з учителем та побудови онтологій для верифікації висловлених тверджень [30]. Аналогічні рішення впроваджують провідні медіа холдинги, серед яких BBC, Reuters та The New York Times, де технології штучного інтелекту застосовуються для автоматизованого моніторингу джерел та внутрішнього фактчекінгу. Вітчизняна практика переважно орієнтована на адаптації відкритих моделей (як-от BERT, RoBERTa) до україномовного контенту, а також на розробці спеціалізованих лінгвістичних правил для аналізу інформаційної достовірності у контексті публікацій [31].

Особливістю систем, орієнтованих на видавничу галузь, є необхідність у збереженні редакційних стандартів стилю, тональності та достовірності джерел. У зв'язку з цим, інтегровані системи достовірності не лише здійснюють оцінку тексту за допомогою алгоритмічних критеріїв, але й забезпечують взаємодію з редактором або аналітиком через інтерфейси з візуалізацією ризиків, балами достовірності та поясненнями рішень [32].

Актуальною тенденцією є впровадження багатфакторного аналізу, що враховує контекстуальні особливості повідомлення, авторитетність джерела, часові відмінності та тематичну релевантність. В окремих системах використано узагальнені індекси достовірності на основі статистичних та евристичних методів. З метою підвищення точності, деякі підходи включають експертне оцінювання як тренувальний корпус або верифікаційний етап, що забезпечує адаптацію моделей до галузевих стандартів журналістики та видавничої політики [24,30].

Таким чином, сучасні системи оцінювання достовірності текстових повідомлень для видавничої та медійної сфери становлять складні гібридні рішення, що поєднують автоматизовані методи з редакторською експертизою. Їх застосування сприяє підвищенню якості інформаційного продукту, зменшенню ризиків поширення дезінформації та підтримці довіри до медіа-ресурсів у цифрову епоху [30,32].

У видавничій сфері, де достовірність інформації має першочергове значення, активно впроваджуються системи, здатні автоматично оцінювати надійність текстових повідомлень. Зважаючи на те, що реальні повідомлення часто містять неповні, суперечливі або контекстуально неоднозначні дані, дедалі більше уваги приділяється використанню нечіткої логіки (fuzzy logic), яка дозволяє моделювати рівень достовірності в умовах невизначеності [40,50].

Системи, побудовані на принципах нечіткої логіки, забезпечують гнучке оцінювання повідомлень на основі лінгвістичних змінних, таких як "низька достовірність", "помірна достовірність", "висока достовірність", замість

жорстких бінарних класифікацій (правда/неправда). Наприклад, розроблені моделі аналізують текстові ознаки — кількість посилань на перевірені джерела, наявність емоційної лексики, частоту тверджень без аргументації — та інтерпретують їх у вигляді нечітких множин. Далі, за допомогою бази правил типу "якщо-то", формується узагальнений рівень довіри до повідомлення [46].

Одним із прикладів таких підходів є розроблення адаптивних експертних систем, де нечітка логіка поєднується з машинним навчанням. У таких системах результати аналізу подаються не лише у вигляді остаточного вердикту, а як шкала достовірності з поясненням причин. [25].

Особливої уваги заслуговують системи, що враховують експертні оцінки при формуванні бази нечітких правил. У таких моделях використано методи агрегування думок спеціалістів (наприклад, редакторів, фактчекерів, лінгвістів), що дозволяє враховувати галузеві специфіки й адаптувати системи до стилістики та стандартів певного видавництва. Такі підходи демонструють високу гнучкість у ситуаціях, коли формальна логіка виявляється недостатньою для об'єктивної оцінки [24].

У контексті україномовного інформаційного простору перспективним є поєднання нечіткої логіки з методами семантичного аналізу, що дозволяє враховувати мовні особливості, контекстуальну поліфонію та латентні смислові структури текстів. Дослідження показують, що системи на основі нечіткої логіки демонструють стабільні результати при оцінюванні повідомлень соціально-політичного характеру, де традиційні алгоритми часто помиляються через надмірну чутливість до поверхневих лексичних індикаторів [40].

Таким чином, застосування нечіткої логіки у видавничих системах верифікації текстових повідомлень забезпечує підвищену адаптивність, прозорість прийняття рішень і точнішу відповідність редакційним стандартам, особливо в умовах семантичної неоднозначності або міждисциплінарного контенту. Поєднання цього підходу з експертною оцінкою та штучним

інтелектом є одним із найперспективніших напрямів розвитку технологій інформаційної перевірки у сфері сучасного видавництва [50].

Таблиця 1.1

Порівняльний аналіз систем оцінювання достовірності текстових повідомлень

Назва системи	Методологія оцінки	Підтримка нечіткої логіки	Джерело знань	Особливості застосування у видавництві
ClaimBuster	Машинне навчання + NLP	Ні	Фактографічна база	Автоматичне виявлення неправдивих тверджень у тексті
Full Fact	Гібридна (AI + експерт)	Ні	Експертна база + API	Внутрішній редакційний фактчекінг у ЗМІ
AdVerif.ai	Семантичний аналіз + ML	Частково	Семантичні графи	Боротьба з дезінформацією в рекламі та медіа
FuzzyTrust (умовна назва)	Нечітка логіка + правила	Так	Експертні правила	Лінгвістична оцінка за шкалою довіри у видавничих системах
HybridCheck-UA (дослідницька модель)	Fuzzy logic + контекстний аналіз	Так	Україномовні джерела + експертні оцінки	Адаптована до локального контенту; використовується в редакціях новин

У таблиці наведено порівняння п'яти систем, орієнтованих на оцінювання достовірності текстів у видавничій сфері, з акцентом на методи, джерела знань та підтримку нечіткої логіки.

ClaimBuster і Full Fact є класичними прикладами систем на основі машинного навчання або гібридного підходу з участю редакторів. Вони демонструють високу ефективність в автоматизованому виявленні фактологічних помилок, але не використовують механізмів нечіткої логіки, що обмежує їхню здатність обробляти повідомлення з невизначеним ступенем достовірності.

AdVerif.ai реалізує часткову підтримку нечітких правил для оцінки контекстуальної достовірності, зокрема в медіа-рекламі, але орієнтована переважно на комерційні тексти.

FuzzyTrust умовно представляє системи, які повністю інтегрують нечітку логіку. Такі системи зазвичай формують шкалу достовірності на основі лінгвістичних змінних та експертних правил, що дає змогу оцінити навіть ті повідомлення, які не мають чітко вираженої фактографічної основи. Їх перевага — висока адаптивність до редакційного процесу.

HybridCheck-UA — приклад дослідницької україномовної системи, що поєднує нечітку логіку, семантичний аналіз і контекстуальні параметри. Вона враховує специфіку локального інформаційного простору, забезпечуючи гнучку та глибоку оцінку публікацій у новинних агентствах.

Таким чином, системи, що інтегрують нечітку логіку, надають видавництвам потужний інструмент для роботи з неоднозначною інформацією, дозволяючи адаптувати механізми перевірки до специфіки стилю, тональності й стандартів конкретного медіа (рис. 1.1).



Рис 1.1. Порівняльна оцінка систем визначення достовірності повідомлень за ключовими критеріями

Їх ефективність особливо проявляється у випадках, де звичайні алгоритми демонструють нестабільні результати через складність інтерпретації мови або контексту.

На гістограмі зображено порівняльну оцінку п'яти систем за чотирма ключовими критеріями:

- ML / NLP – рівень використання машинного навчання та обробки природної мови;
- Fuzzy Logic – ступінь інтеграції нечіткої логіки;
- Expert Input – роль експертних знань у роботі системи;
- Publishing Adaptability – пристосованість системи до потреб видавничої сфери.

Діаграма демонструє, що FuzzyTrust і HybridCheck-UA мають найвищу підтримку нечіткої логіки та найкращу адаптацію до видавничого середовища, тоді як ClaimBuster та Full Fact більше зосереджені на автоматичній обробці фактів за допомогою машинного навчання.

Аналіз сучасних систем визначення достовірності текстових повідомлень засвідчив значну різноманітність підходів, які поєднують методи машинного навчання, семантичного аналізу, експертного оцінювання та нечіткої логіки. Системи, що застосовуються у видавничій сфері, мають враховувати специфіку журналістського тексту — його багатозначність, жанрову різноманітність і соціальний вплив. Найефективнішими виявлено гібридні моделі, що поєднують автоматизовану обробку з редакторською експертизою та контекстною валідацією повідомлень [32].

У медійній сфері системи перевірки достовірності відіграють дедалі важливішу роль у боротьбі з дезінформацією та забезпеченні етичної журналістики. Вони інтегровані у виробничі цикли редакцій та допомагають формувати прозорий і безпечний інформаційний простір. Особливо актуальними є ті інструменти, які дозволяють працювати з текстами оціночного характеру, прогнозами та інтерпретаціями, забезпечуючи при цьому аналітичну гнучкість і

зрозумілість результатів [40].

Правові аспекти використання таких систем є невід'ємною складовою їхнього впровадження [33-35]. Міжнародні та національні нормативно-правові акти формують правову базу, що регулює обіг інформації, захист персональних даних, прозорість алгоритмічного аналізу та відповідальність за поширення недостовірної інформації. Успішне застосування технологій у медіа залежить не лише від їхньої технічної ефективності, а й від дотримання прав людини, етичних стандартів журналістики та принципів прозорості.

Інтеграція систем автоматичного оцінювання достовірності текстових повідомлень у видавничу та медійну діяльність супроводжується важливими правовими викликами. Вони стосуються не лише забезпечення прав громадян на достовірну інформацію, а й регулювання роботи інтелектуальних алгоритмів, збереження авторських прав, етичного використання штучного інтелекту та дотримання принципів свободи слова.

На національному рівні правове регулювання в інформаційній сфері здійснюється відповідно до таких актів.

Закон України «Про основні засади забезпечення кібербезпеки України» [33], який встановлює право громадян на отримання об'єктивної, повної, достовірної інформації. У сфері медіа та видавничої справи діє «Кодекс України про адміністративні правопорушення щодо відповідальності за поширення завідомо неправдивих відомостей» [34]. У період воєнного стану значення набувають нормативи, що стосуються інформаційної безпеки держави – Закон України «Про заборону пропаганди російського нацистського тоталітарного режиму, збройної агресії Російської Федерації як держави-терориста проти України, символіки воєнного вторгнення російського нацистського тоталітарного режиму в Україну» [35]. що, у свою чергу, підвищує вимоги до достовірності журналістського контенту.

Правове забезпечення функціонування систем оцінювання достовірності текстових повідомлень у медійній та видавничій сферах є ключовою умовою

їхнього легітимного впровадження та етичного використання. Проаналізовані нормативно-правові акти свідчать про існування як національного, так і міжнародного правового поля, що регламентує обіг інформації, відповідальність за її недостовірність, правила обробки персональних даних та обмеження щодо використання штучного інтелекту.

Отже, впровадження систем оцінювання достовірності повинно здійснюватися у межах чітко окресленого правового середовища, з дотриманням принципів прозорості, справедливості, недискримінації та підзвітності, що забезпечує як правову безпеку користувачів, так і підвищення якості публічного інформаційного простору.

Загалом аналіз зазначених систем свідчить, що автоматизовані засоби перевірки інформації мають значний потенціал для підтримки користувачів у процесі ідентифікації недостовірних повідомлень. Їх здатність швидко обробляти великі обсяги даних, виявляти ключові ознаки дезінформації та формувати попередні висновки щодо достовірності контенту робить їх цінними у сфері інформаційної безпеки. Водночас наявні рішення демонструють потребу в подальшому вдосконаленні, зокрема — у підвищенні точності алгоритмічного виведення, що має ґрунтуватися на глибшому аналізі контексту.

1.4. Засоби факторного аналізу для дослідження проблематики оцінювання рівня достовірності текстових повідомлень

Описані вище методи і програмні засоби забезпечують встановлення достовірності текстових повідомлень фактично після їх випуску, тобто користувач оцінить правдивість отриманої інформації вже після її отримання. З іншого боку, для прийняття адекватного висновку пересічний користувач повинен мати комп'ютер та володіти відповідним програмним забезпеченням, що визначить характер та істинність повідомлення. Зрозуміло, що подібний контроль є затратним у матеріальному і часовому вимірі, часто потребує належних приладів та кваліфікованих виконавців, хоча при наявності

відповідних передумов забезпечує достовірність рішень, прийнятих на основі отриманих текстових даних.

Додатково проблема полягає у відсутності правил та стандартів щодо власне процесу творення, передавання та розповсюдження подібної інформаційної продукції. В останні роки вона виступає не лише засобом реклами (не завжди порядної), але й бізнесових та політичних маніпуляцій, кіберзалякувань, які мають істотний вплив на великі суспільні групи [37-39].

Завдання пропонованого дослідження полягає у розробленні додаткових засобів (теоретичних і прикладних), які у сукупності з існуючими методами уможливили б створення ефективного механізму прогностичного оцінювання достовірності текстових повідомлень до моменту їх отримання користувачем. Основою пропонованого підходу могли б стати такі чинники, як оцінювання особи автора повідомлень та джерела інформації, соціальна довіра до описаних фактів, критичне сприйняття з метою спростування фейкових новин і т. п.

Зрозуміло, що перераховані вище та подібні до них фактори крім словесного опису не можуть бути задані числовими характеристиками, що унеможлиблює їх дослідження традиційними математичними методами. Виходом з цієї ситуації є використання теорії нечітких множин [40-46], які використовуються для моделювання ситуацій, де звичні (чіткі) множини неефективні через нечіткість або неоднозначність реальних даних. Узагальненою величиною при цьому виступає лінгвістична змінна, значення якої задається словесно описовим способом. Кожне з цих значень відображено нечіткою множиною, або терм-множиною. Значення лінгвістичних змінних задано функціями належності, які показують, наскільки кожне значення змінної відповідає певному терміну. За їхнього застосування лінгвістичні характеристики трансформовано у числові представлення, що дає змогу здійснювати комп'ютерну обробку моделей, спрямованих на оцінювання достовірності текстової інформації. Вказаний підхід уможливить опрацьовувати інформацію, яка є нечіткою, неточною або неоднозначною, забезпечуючи при

цьому прийняття рішень на основі наближених даних.

У нашому дослідженні лінгвістичні змінні, що стосуються завдань із соціальним спрямуванням, розглянуто як фактори, які впливають на прогнозоване визначення рівня достовірності текстових інформаційних повідомлень. Завдання вирішено на інформаційному рівні, оскільки прогноз побудовано на основі розроблених моделей, у яких ключову роль відіграє порівняння факторів через їх попарне зіставлення та визначення пріоритету кожного з них у процесі дослідження. Цей підхід є універсальним, оскільки він дає змогу застосовувати єдину методологію як до числових величин, так і до менш формалізованих вимог, характеристик, описаних природною мовою. Важливою перевагою використання нечіткої логіки є її здатність ефективно працювати з неточними, неоднозначними або частково визначеними даними, що часто зустрічаються в реальних ситуаціях. Це робить нечітку логіку корисною для широкого спектра завдань і сфер діяльності [46, 47, 50].

Оскільки апріорі вплив різних чинників на перебіг процесу є неоднаковим, доцільно для реалізації озвученого напрямку досліджень використати методи та засоби, що забезпечать формалізоване відображення зв'язків між факторами з використанням графів семантичних мереж – вихідної передумови визначення переваг факторів за пріоритетністю дії на процес та отримання відповідних багаторівневих моделей впливу виокремлених факторів на показник достовірності текстових даних. Важливою є проблема числового вираження міри впливу фактора на процес, що вирішується через попарні порівняння факторів з додержанням принципу числової або кардинальної погодженості за рівнем пріоритетності [46]. У результаті його застосування отримано умовні вагові значення факторів, що стають базовою інформаційною основою для подальших досліджень.

Для кращого розуміння суті перетворень над факторами, що характеризують достовірність текстових інформаційних повідомлень, наведемо теоретичні викладки, які стосуються перетворень над лінгвістичними змінними

[46-48,50]. Допустимо існування факторів $x_1, \dots, x_n$ деякого процесу, оформлених у вигляді семантичної мережі – орієнтованого графа, у вершинах якого розміщено власне фактори, а ребра визначають лінгвістичну суть зв'язку між ними. Опрацювання мережі за методом аналізу ієрархій та методом ранжування приводить до отримання багаторівневої моделі, що відтворює ступені впливу факторів на процес. Крім стратифікованих рівнів визначають умовні ваги $p_1, \dots, p_n$ факторів через експертне порівняння їх між собою. Переваги фактора x_i по відношенню до фактора x_j позначимо величинами a_{ij} , сформувавши з них обернено симетричну матрицю $A = (a_{ij})$, в якій згідно визначення $a_{ij} = 1/a_{ji}$. Отримана таким чином матриця вважається узгодженою, якщо остання рівність справедлива для всіх порівнянь умовних ваг факторів між собою. Для узгодженої обернено симетричної матриці очевидним є таке співвідношення між вагами факторів:

$$a_{ij} = \frac{p_i}{p_j} \text{ для всіх } i, j = 1, 2, \dots, n. \quad (1.1)$$

Загалом матричне рівняння $Ax = y$ слугує аналогом системи рівнянь

$$\sum_{j=1}^n a_{ij} x_j = y_i; i = 1, 2, \dots, n, \quad (1.2)$$

яка з врахуванням співвідношення (1.1) може бути подана таким чином:

$$\sum_{j=1}^n a_{ij} p_j = np_i; i = 1, 2, \dots, n. \quad (1.3)$$

Остання рівність у векторному форматі матиме такий вигляд:

$$Ap = np. \quad (1.4)$$

Величина p у виразі (1.4) означає власний вектор обернено симетричної матриці попарних порівнянь факторів A з власним значенням n .

Згідно з основоположною концепцією теорії факторного аналізу [46] за умови наявності факторів, що вважаються «рушіями» довільного процесу, впливи та взаємозв'язки між ними визначаються на підставі експертних суджень, що безумовно вносить елементи суб'єктивізму у процес дослідження. Початкові твердження експертів про наявність чи відсутність зв'язку (впливу) між факторами через бінарні значення логічної змінної («так» – зв'язок існує, «ні» – його немає) не забезпечує розрахунок точного значення змінної a_{ij} за формулою (1.1). У результаті таких порівнянь можна формують бінарну матрицю досяжності, опрацювання якої за методом аналізу ієрархій приведе (стосовно напряму дисертаційного дослідження) до синтезу багаторівневої моделі переваг факторів стосовно впливу на достовірність текстових даних.

Продовження дослідження потребує застосування методів, які на підставі стратифікованої моделі рангів факторів забезпечать достовірність остаточних результатів у формі числових вагових значень чинників процесу, розрахованих в умовних одиницях. Реалізація такого підходу можлива за умови використання положень з теорії матриць [46, 51], суть яких наведемо нижче.

Отже, при наявності чисел $\lambda_1, \dots, \lambda_n$, що задовольняють рівняння $Ax = \lambda x$, тобто є власними значеннями матриці A , причому $a_{ii} = 1$ ($i = 1, \dots, n$), то для сукупності значень λ_i матимемо таку умову: $\sum_{i=1}^n \lambda_i = n$. Остання рівність при наявності формули (1.2) означає, що у випадку узгодженості експертних суджень тільки одне максимальне значення власного вектора матриці A дорівнюватиме n , тобто $\lambda_{\max} = n$ (або близьке до нього при неістотній зміні елементів a_{ij}). Решта

значень нульові. Цю величину застосовано для встановлення міри узгодженості експертних суджень щодо попарних порівнянь факторів, виражених обернено симетричною матрицею. Таким чином, міра відхилення $\lambda_{\max}$ від n може служити критерієм узгодженості експертних суджень стосовно ваг факторів стосовно рівня їх розміщення в багаторівневій моделі. Він називається індексом узгодженості і розраховується за формулою

$$IU = \frac{\lambda_{\max} - n}{n - 1}. \quad (1.5)$$

Задовільний результат порівняння факторів при $IU \leq 0,1$.

З урахуванням вище наведеного та на підставі багаторівневої моделі пріоритетів факторів встановлюємо експертним способом попередні вагові значення факторів, що слугують основою для побудови матриці попарних порівнянь за шкалою відносної важливості об'єктів [48]. Опрацювання матриці за програмою, суть якої буде розкрита в наступному розділі, приводить до отримання вектор пріоритетів, нормалізовані компоненти якого вважаються оптимізованими ваговими пріоритетами факторів оцінювання достовірності текстових даних.

Отримані вагові значення факторів стають основою для проектування альтернативних та розрахунку оптимального варіантів виконання процесу визначення рівня достовірності даних. При цьому для розв'язання вказаного завдання доцільно враховувати положення, сформульовані нижче.

Положення 1.1. Альтернативні варіанти процесу встановлення рівня достовірності текстових даних формуються з набору факторів, що задовольняють принципу Парето, тобто пріоритетність впливу яких на правдивість даних є визначальною [48].

Суть положення означає, що для виокремленої та упорядкованої за перевагами впливу на процес множини факторів $\{x_1, x_2, \dots, x_n\}$, множину Парето складатиме

частина пріоритетних факторів $\{x_1, x_2, \dots, x_{n_1}\}$, де $n_1 < n$.

Положення 1.2. Для формування альтернативних варіантів використовуються вагові пріоритети факторів та встановлюються часткові міри оцінювання альтернативи стосовно впливу факторів множини Парето на виконання процесу за даною альтернативою, тобто:

$$AV = F[x_i(w_i), p_{ij}], \quad (1.6)$$

де: x_i – фактори множини Парето; w_i – ваговий показник пріоритетності i -го фактора; p_{ij} – величина оцінювання j -ї альтернативи за мірами важливості i -го фактора, виражена у відсотках.

На завершальному етапі відзначимо важливість та особливість формування бази знань. Вона містить:

1. Лінгвістичні змінні – опис параметрів системи через нечіткі множини (наприклад, «низький», «середній», «високий»).
2. Функції належності – математичні залежності, які визначають, наскільки елемент належить до певної нечіткої множини.
3. Нечіткі правила – сукупність умовно-логічних залежностей у форматі «Якщо-Тоді».

Використання засобів нечіткої логіки на завершальному етапі реалізації інформаційного підходу щодо встановлення рівня достовірності текстових повідомлень передбачає такі кроки [46].

1. Розроблення методу логічного виведення як вихідної інформаційної бази, багаторівнева структура якої відтворює ієрархію факторів та відповідних їм лінгвістичних термів значень лінгвістичних змінних.
2. Візуалізація за нормованими значеннями функцій належності для лінгвістичних змінних і відповідних їм лінгвістичних термів відносно квантів поділу інтервалів значень універсальної терм-множини.
3. Побудова нечітких матриць знань – передумови забезпечення залежностей між вхідними та вихідними даними на основі логічних правил.

4. Формування та розрахунок нечітких логічних рівнянь, що ідентифікують зв'язки між функціями належності та лінгвістичними термами.
5. Вибірка даних для розрахунку прогнозного значення показника рівня достовірності текстових інформаційних повідомлень.

На завершення наведемо розроблену багаторівневу модель виконання дисертаційного дослідження (рис. 1.2).



Рис. 1.2. Модель дослідження процесу оцінювання достовірності текстових інформаційних повідомлень

У процесі реалізації моделі, спрямованої на аналіз й оцінювання достовірності текстових інформаційних повідомлень, досягнуто суттєвих результатів. Поєднання традиційних методів з новітніми підходами дало змогу сформувати інтегровану інформаційну технологію, здатну ефективно виявляти ознаки недостовірності та забезпечувати прогнозування рівня достовірності за допомогою інструментарію нечіткої логіки. Такий підхід забезпечує підвищення точності та оперативності в оцінюванні інформації, що набуває особливого значення в умовах поширення фейкових новин і маніпулятивного контенту. Залучення факторів, які впливають на формування довіри до джерел інформації, суттєво підвищує достовірність результатів і сприяє глибшому розумінню надійності інформаційних повідомлень.

У межах практичного втілення моделі передбачено розроблення й удосконалення методів, інтегрованих в інформаційну технологію, яка базується на класифікаційних групах лінгвістичних змінних — організаційній, технічній і користувацькій. Ці групи відображають етапи реалізації процедур формування показника достовірності текстових повідомлень. Запропонована модель має потенціал забезпечити об'єктивне та високошвидкісне прогнозне оцінювання новинного контенту.

1.6. Завдання дисертаційної роботи

Для вирішення проблематики встановлення достовірності текстових інформаційних повідомлень необхідно розв'язати такі завдання:

- виконати аналіз методів, засобів та технологій оцінювання рівня достовірності інформаційних повідомлень;
- розробити моделі визначення оптимального варіанту процесу формування інтегрального показника оцінювання достовірності текстових інформаційних повідомлень;
- розробити метод логічного виведення та сформувати нечіткі матриці знань та нечіткі логічні рівняння для лінгвістичних рівнів моделі;

- удосконалити метод формування показника достовірності текстових повідомлень та визначення вірогідності контенту отриманих текстових повідомлень, у якому реалізовано фазифікацію як процес переходу від числових даних до нечітких оцінок на основі лінгвістичних змінних;

- виконати дефазифікацію нечітких множини, реалізовану на основі нечіткого логічного виведення, що забезпечить отримання для заданих варіантів вихідних даних показника прогнозованого рівня достовірності текстових повідомлень;

- розробити інформаційну технологію прогнозного оцінювання достовірності текстових інформаційних повідомлень;

Висновки до розділу 1

1. Виконано аналіз методів, засобів та технологій оцінювання рівня достовірності інформаційних повідомлень. Відзначено, що вивчення і дослідження вказаної тематики вимагає комплексного підходу для розв'язання складних проблем, пов'язаних з достовірністю інформації в сучасному інформаційному середовищі

2. Описано можливі напрями та варіанти реалізації результатів досліджень, що стосуються оцінювання достовірності текстових повідомлень та різного роду інформаційних повідомлень.

3. Наведено перелік теоретичних і прикладних ресурсів факторного аналізу, використаних для розроблення моделей і засобів процесу дослідження проблематики оцінювання достовірності текстових даних. До них віднесено:

- семантичні мережі – для формалізованого відтворення та опису зв'язків між факторами впливу на рівень достовірності текстових повідомлень;

- метод аналізу ієрархій та метод ранжування – для розроблення багаторівневої моделі важливості факторів;

- метод попарних порівнянь та шкалу відносної важливості об'єктів – для визначення вагових переваг факторів та синтезування оптимізованої моделі

пріоритетного впливу факторів рівень достовірності текстових повідомлень;

- метод лінійного згортання критеріїв – для проектування альтернативних та розрахунку оптимального варіантів визначення достовірності повідомлень;

- метод нечіткого логічного виведення з урахуванням класифікаційних груп лінгвістичних змінних – організаційної, технічної та користувацької – який відображає поетапну реалізацію процедур формування показника достовірності текстових інформаційних повідомлень.

- функції належності лінгвістичних змінних, нечіткі матриці знань та нечіткі логічні рівняння – передумова розрахунку показника прогнозованого рівня достовірності текстових інформаційних повідомлень;

- структурно-функціональну модель інформаційної технології визначення рівня достовірності текстових повідомлень, що апріорі обумовлює правдивість інформаційних даних;

4. Розроблено багаторівневу модель виконання дисертаційного дослідження, орієнтовану на оцінювання достовірності текстових інформаційних повідомлень. Модель базується на аналізі, який містить синтаксичний, семантичний та прагматичний рівні, що забезпечить високий рівень точності оцінювання отримуваних даних.

5. Сформульовано основні завдання дисертаційної роботи.

РОЗДІЛ 2. МОДЕЛІ ФАКТОРІВ ВПЛИВУ НА ДОСТОВІРНІСТЬ ТЕКСТОВИХ ІНФОРМАЦІЙНИХ ПОВІДОМЛЕНЬ

2.1. Фактори процесу формування рівня достовірності текстових повідомлень

Аналіз літературних джерел підтвердив наявність різнопланових чинників, що мають суттєвий вплив на достовірність інформації, переважаючи частку якої складають текстові повідомлення. Крім наведених робіт з даної тематики попереднього розділу, доцільно звернути увагу на декілька публікацій, суттєво близьких до завдань даного розділу [36-38]. Вони обумовили більш обґрунтований вибір факторів процесу формування рівня достовірності інформаційних повідомлень.

Додатковою передумовою вважаємо наступні міркування. У видавничій сфері перелік зазначених чинників достовірності текстових даних формується експертами редакційного процесу, журналістської етики та стандартів професійної комунікації. Фахівці цієї галузі виходять з того, що видавничий продукт безпосередньо впливає на суспільну думку, а отже, вимоги до його достовірності є вкрай високими. При цьому вибір основних чинників достовірності інформації здійснено із застосуванням методу узгодженості експертних оцінок, зокрема статистики Каппа Кохена [39]. Цей підхід дозволяє кількісно оцінити рівень згоди між незалежними експертами щодо важливості окремих показників, що забезпечує обґрунтованість і наукову валідацію отриманих результатів.

Під час експерименту групі фахівців у сфері видавничої діяльності було запропоновано визначити ступінь значущості впливових факторів достовірності інформації, яка публікується. Узгодженість їхніх відповідей проаналізовано за допомогою коефіцієнта Каппа, який враховує випадкову ймовірність збігу оцінок. Високі значення цього показника засвідчили, що експерти демонструють стійку та статистично значущу згоду щодо визначених критеріїв.

За результатами експертного оцінювання та враховуючи роботи [32-36], до

переліку важливих факторів достовірності інформації, релевантних для видавничої сфери, увійшов деякий узагальнений набір чинників. Наведемо їх назву, мнемонічне позначення та коротку функціональну характеристику.

Професіоналізм автора (ПА-1) у видавничій сфері є ключовим критерієм, оскільки саме компетентність журналіста, редактора чи автора дозволяє гарантувати належний рівень аналітичності, точності та обґрунтованості інформації. *Об'єктивність подачі матеріалу (ОА-2)* є основою редакційної політики авторитетних видань, які прагнуть зберігати довіру аудиторії та уникати маніпулятивних чи суб'єктивних висловлювань. *Інформативність контенту (ІК-3)* є важливою складовою якісного видавничого продукту, оскільки читачі очікують надання повної та структурованої інформації для формування власного обґрунтованого погляду. *Джерело інформації (ДІ-4)* має особливе значення для видавничої сфери, адже посилання на офіційні документи, експертні коментарі чи перевірені бази даних є стандартом для професійних редакцій. *Перевірка фактів (ПФ-5)* активно використовується у провідних видавництвах через наявність внутрішніх процедур фактчекінгу, які дозволяють мінімізувати ймовірність помилок чи свідомої дезінформації. *Багаторазове публікування (БП-6)* певної інформації в різних незалежних джерелах часто є критерієм, за яким редакції оцінюють доцільність її використання у власних матеріалах. Вагоме значення у видавничій сфері має здатність інформації витримувати публічне *спростування і критику (СК-7)*, оскільки матеріали, які не піддаються суттєвим запереченням після публікації, зміцнюють авторитет видання. *Соціальна довіра (СД-8)* до джерела, у даному випадку до конкретного видавництва або редакції, є визначальним чинником довгострокового збереження репутації та лояльності аудиторії. *Актуальність (АК-9)* інформації повинна бути беззаперечною, оскільки застаріла інформація може втратити свою достовірність. Необхідно перевіряти дату публікації та оновлення інформації. *Мова та стиль написання (МС-10)*. Професійна та грамотна мова може свідчити про серйозність та дбайливість автора щодо

надійності інформації. *Методологія дослідження (МД-11)* має сенс, якщо текст повідомлення містить результати досліджень, оскільки це може впливати на його надійність. *Технічні аспекти (ТА-12)* обумовлюють врахування безпекових аспектів, таких як захищеність сайту, наявність сертифікатів безпеки SSL (Secure Sockets Layer – безпечний протокол забезпечення з'єднання), а також перевірку на технічні помилки чи редагування. Зараз SSL замінений протоколом TLS (Transport Layer Security – засіб забезпечення безпеки транспортного рівня).

Використання експертного оцінювання обумовило одержання кількісної оцінки коефіцієнтів важливості кожного з факторів, що формують множину значень чинників впливу на достовірність текстових повідомлень. Візуалізація результатів експертного оцінювання наведена на рис. 2.1.

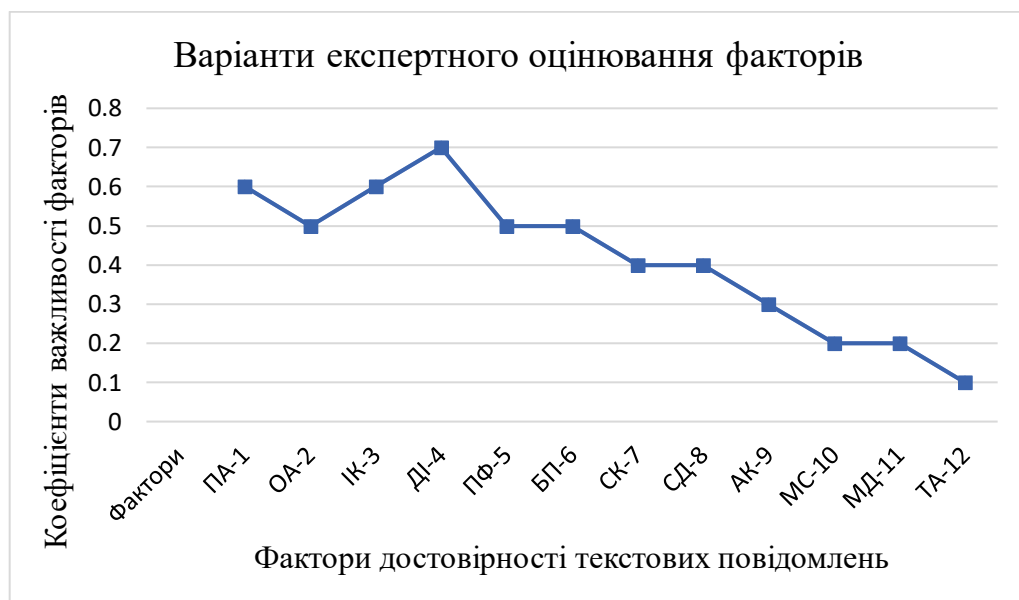


Рис. 2.1. Варіанти експертного оцінювання вагомості факторів достовірності текстових інформаційних повідомлень

В цьому контексті наведемо одне з важливих положень, суть якого має пряме відношення до обґрунтування прийнятої для дослідження концепції [46].

У контексті виробничих, соціальних і технологічних систем довільного характеру їх функціонування зумовлене сукупністю факторів, що впливають на кінцеві результати. Серед усього спектра таких чинників вирізнено підмножину, яка відіграє ключову роль у формуванні якісних характеристик відповідних процесів. З іншого боку, часто не менш важливими є результати, обумовлені

процедурами формування та визначення рівня достовірності інформації про вказані процеси. Формалізований запис вказаного твердження може мати таке вираження.

Нехай $R = \{r_1, r_2, \dots, r_m\}$ – певна множина деяких процесів; $M = \{x_{1_m}, x_{2_m}, \dots, x_{n_m}\}$ – множина впливових факторів на ступінь достовірності даних стосовно m -го процесу, а n_m – число факторів цього процесу. Крім того, процес описано певною функцією $A(x)$, що визначає не тільки вірогідність та адекватність результатів виконання процесу, але й достовірність даних про нього. Узагальнюючи, можна прийняти такий формалізований запис

$$A(x) \equiv \bigcup_{j=1}^n \omega(x_{j_k}), \quad (k=1, 2, \dots, m), \quad (2.1)$$

де: $A(x)$ – числове вираження функції достовірності m -го процесу; $\omega(x_{j_k})$ – числова міра достовірності, привнесеної j -м чинником у k -й процес.

Беручи до уваги співвідношення (2.1), викладене положення представлено наступним чином: у множині можливих процесів обов'язково існує хоча б один, для якого всі залучені фактори задовольняють умову (2.1), формалізовану у вигляді такого виразу

$$(\exists r) (\forall x) A(x); r \in R; x \in M. \quad (2.2)$$

Виокремимо із загального набору факторів деяку підмножину R , що стане у майбутньому основою формування рівня достовірності текстових повідомлень. Виразимо залежність показника через змінні другого порядку:

$$R = F_R(A, Z, S), \quad (2.3)$$

де: A – змінна, яка позначає підмножину факторів, що свідчать про кваліфікацію автора повідомлення; Z – змінна ідентифікації факторів організаційних заходів; S – змінна факторів суспільної лояльності.

Розкриємо вказані підмножини, прикріпивши до них відповідні фактори.

$$A = F_A(a_1, a_2, a_3), \quad (2.4)$$

де: a_1 – професіоналізм автора; a_2 – об’єктивність автора; a_3 – інформативність контенту повідомлення.

$$Z = F_Z(z_1, z_2, z_3), \quad (2.5)$$

де: z_1 – джерело інформації; z_2 – перевірка фактів; z_3 – багаторазове публікування (перевірка через пошук багаторазового публікування).

$$S = F_S(s_1, s_2), \quad (2.6)$$

де: s_1 – спростування і критика; s_2 – соціальна довіра.

Подальші дії виконаємо над об’єднаною множиною факторів, елементи якої для спрощення матимуть однотипне позначення:

$$X = \{x_1, x_2, x_3, x_4, x_5, x_6, x_7, x_8\}. \quad (2.7)$$

Для початку відтворимо зв’язки між факторами множини (2.7) за допомогою орієнтованих графів, а саме їх різновиду – семантичних мереж.

2.2. Формалізоване відтворення зв’язків між факторами засобами семантичних мереж та мови предикатів

Розшифруємо множину факторів (2.7), доповнивши їх математичне трактування семантичною сутністю, зрозумілою для відображення.

$$X = \left\{ \begin{array}{l} x_1 - \text{професіоналізм автора}; \\ x_2 - \text{об'єктивність автора}; \\ x_3 - \text{інформативність контенту}; \\ x_4 - \text{джерело інформації}; \\ x_5 - \text{перевірка фактів}; \\ x_6 - \text{багаторазове публікування}; \\ x_7 - \text{спростування і критика}; \\ x_8 - \text{соціальна довіра}. \end{array} \right\} \quad (2.8)$$

На основі множини (2.8) та результатів експертного опитування щодо можливих відношень між чинниками впливу на достовірність даних, запроектовано графічну модель зв'язків між факторами (рис. 2.1) [53,54].

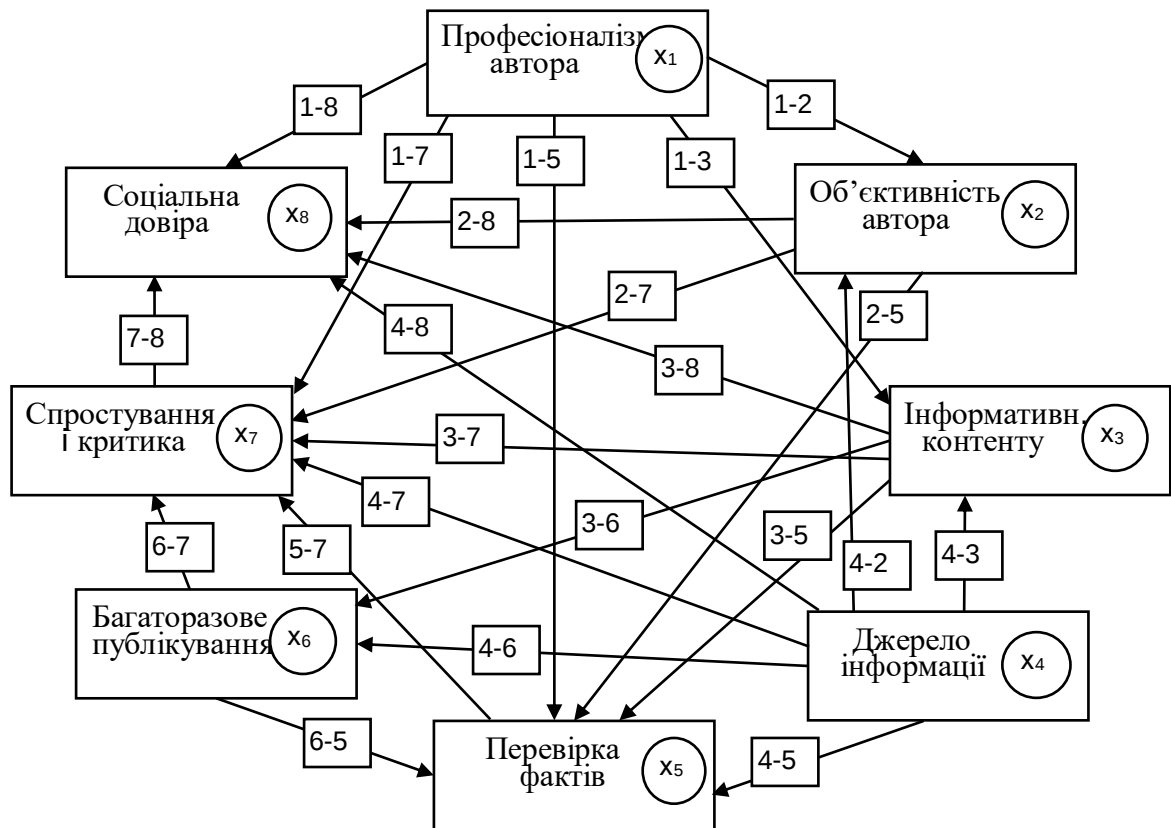


Рис. 2.1. Семантична мережа факторів достовірності текстових повідомлень

Вершини орієнтованого графа, що відтворює семантичну мережу, позначено факторами множини $X = \{x_1, x_2, \dots, x_8\}$, зв'язки між вершинам (x_i, x_j) – дугами, на яких для кращого орієнтування додатково вказано джерело та

приймач впливу.

З огляду на семантичну мережу рис. 2.1 здійснено формалізацію відношень між факторами використовуючи мову предикатів, що передбачає прості предикати та логічні умови: $\wedge$ – «і»; $\vee$ – «або»; $\leftarrow$ – «якщо»; $\forall$ – знаки спільності та $\exists$ – існування [18].

Опис зв'язків між факторами забезпечили такі предикати [18]:

- предикати впливу – визначає, формує, обумовлює, стає основою, передбачає, враховує, діє;
- предикати залежності – визначається, формується, обумовлюється, ґрунтується, передбачається, враховується, отримує.

З урахуванням наведених передумов представлено формалізоване вираження відношень між факторами, використовуючи графічну семантичну мережу рис. 2.1

$(\forall x_i)[\exists(x_1, \text{професіоналізм автора}) \leftarrow \text{обумовлює}(x_1, x_2) \wedge \text{визначає}(x_1, x_3) \wedge \text{враховує}(x_1, x_5) \wedge \text{передбачає}(x_1, x_7) \wedge \text{формує}(x_1, x_8)];$

$(\forall x_i)[\exists(x_2, \text{об'єктивність автора}) \leftarrow \text{стає основою}(x_2, x_5) \wedge \text{діє}(x_2, x_7) \wedge \text{обумовлює}(x_2, x_8) \wedge \text{обумовлюється}(x_2, x_1) \wedge \text{отримує}(x_2, x_4)];$

$(\forall x_i)[\exists(x_3, \text{інформативність контенту}) \leftarrow \text{визначає}(x_3, x_5) \wedge \text{передбачає}(x_3, x_6) \wedge \text{діє}(x_3, x_7) \wedge \text{стає основою}(x_3, x_8) \wedge \text{визначається}(x_3, x_1) \wedge \text{обумовлюється}(x_3, x_4)];$

$(\forall x_i)[\exists(x_4, \text{джерело інформації}) \leftarrow \text{діє}(x_4, x_2) \wedge \text{обумовлює}(x_4, x_3) \wedge \text{враховує}(x_4, x_5) \wedge \text{передбачає}(x_4, x_6) \wedge \text{визначає}(x_4, x_7) \wedge \text{формує}(x_4, x_8)];$

$(\forall x_i)[\exists(x_5, \text{перевірка фактів}) \leftarrow \text{обумовлює}(x_5, x_7) \wedge \text{враховується}(x_5, x_1) \wedge \text{ґрунтується}(x_5, x_2) \wedge \text{визначається}(x_5, x_3) \wedge \text{враховується}(x_5, x_4) \wedge \text{передбачається}(x_5, x_6)];$

$(\forall x_i)[\exists(x_6, \text{багаторазове публікування}) \leftarrow \text{передбачає}(x_6, x_5) \wedge \text{визначає}$

$\wedge (x_6, x_7) \wedge \text{передбачається } (x_6, x_3) \wedge \text{передбачається } (x_6, x_4)];$

$(\forall x_i)[\exists(x_7, \text{спростування і критика}) \leftarrow \text{стає основою } (x_7, x_8) \wedge \text{передбачається } (x_7, x_1) \wedge \text{отримує } (x_7, x_2) \wedge \text{отримує } (x_7, x_3) \wedge \text{визначається } (x_7, x_4) \wedge \text{обумовлюється } (x_7, x_5) \wedge \text{визначається } (x_7, x_6)];$

$(\forall x_i)[\exists(x_8, \text{соціальна довіра}) \leftarrow \text{формується } (x_8, x_1) \wedge \text{обумовлюється } (x_8, x_2) \wedge \text{ґрунтується } (x_8, x_3) \wedge \text{формується } (x_8, x_4) \wedge \text{ґрунтується } (x_8, x_7)].$

Реалізація формалізованого опису семантичної мережі забезпечує упорядкування факторів за пріоритетністю впливу на достовірність текстових повідомлень. Додатково доцільно взяти до уваги певні міркування, смисл яких впливає з логічного трактування певних положень факторного аналізу, застосованих у дослідженні видавничо-поліграфічних процесів [46]. Сформулюємо їх у вигляді поданих нижче означень.

Означення 1. Характер дії чинників (факторів) на процес, у якому вони задіяні, визначено їх перевагою, яку доцільно відтворити деяким умовним ваговим коефіцієнтом. Отриманий таким чином ранг фактора обумовлює рівень його важливості у багаторівневій моделі чинників процесу. Суттєвим при цьому вважається твердження, що множина факторів містить хоча б один найбільш визначальний чинник. Скорочено озвучений висновок можна сформулювати таким чином: для множини $W = \{w_{1_m}, w_{2_m}, \dots, w_{n_m}\}$ ваг факторів певного процесу при виконанні умови $B(w) \equiv \max\{w_{1_m}, w_{2_m}, \dots, w_{n_m}\}$ справедливим буде таке трактування:

$$(\forall r)(\exists w) B(w); r \in R; w \in W. \quad (2.9)$$

Означення 2. У множині n факторів деякого процесу, упорядкованих за їх умовними ваговими значеннями, відсутні однакові за ступенем впливу на процес чинники.

Задеклароване твердження особливо дієве у суспільних процесах, де його актуальність і достовірність звичайно підтверджується результатами діяльності певних особистостей, що виступають «учасниками», або чинниками процесу.

Таким чином, за відсутності рівних за пріоритетами факторів, що виражено умовою $D(w) \equiv w_j > w_{j+1}$ для $(j = 1, 2, \dots, n-1)$, матимемо:

$$(\forall w) D(w); w \in W. \quad (2.10)$$

З урахуванням статичних вихідних понять початкового етапу дослідження переходимо до моделей, що визначають динаміку впливу факторів на достовірність текстових інформаційних повідомлень.

2.3. Модель ієрархічного впливу факторів на формування показника достовірності текстових інформаційних повідомлень

Модель ієрархічного впливу факторів на формування показника достовірності текстових повідомлень дозволяє систематизувати інформацію про ключові аспекти, які впливають на достовірність текстів. Модель орієнтована на аналіз різних чинників, такі як джерело даних, методи збору і обробки, контекст і тематика тексту, кваліфікація автора, та визначає їхню вагомість у формуванні остаточної оцінки достовірності. Вона дозволяє не лише виявляти потенційні джерела помилок або спотворень, але і розробляти стратегії для покращення якості текстової інформації, що є важливим в контексті сучасних вимог до надійності даних.

Дослідження в цій області спрямовано на розроблення моделей і алгоритмів, які враховують дію певних факторів на рівень достовірності окремих текстів або наборів даних. Остаточно уможлиблюється науково обґрунтоване досягнення більш точних і надійних висновків з використанням текстової інформації в різних сферах діяльності, що забезпечує підґрунтя для прийняття

обґрунтованих рішень та стратегій.

Для визначення ступенів пріоритетності факторів, використаних для дослідження процесу визначення достовірності текстових повідомлень, застосовано метод математичного моделювання ієрархій [46,55,56]. Метод є потужним інструментом прийняття рішень, оскільки забезпечує структурування та аналіз складних проблем з урахуванням дії різнорідних чинників.

За алгоритмом реалізації вказаного підходу побудовано квадратну бінарну матрицю досяжності B (таблиця 2.1), що відтворює залежності між факторами семантичної мережі. Елементи матриці визначено за таким правилом:

$$b_{ij} = \begin{cases} 1, \text{ якщо з вершини } i \text{ можна потрапити у вершину } j; \\ 0 \text{ в іншому випадку.} \end{cases} \quad (2.11)$$

З виразу (2.11) автоматично випливає, що діагональні елементи такої матриці рівні одиниці. Можна також стверджувати, що вершина x_j ($j=1, 2, \dots, 8$) буде вважатися досяжною відносно вершини x_i ($i=1, 2, \dots, 8$), якщо з останньої «прямим шляхом» можна потрапити в x_j . Справедливим буде і зворотне твердження: вершину x_i вважатимемо попередницею вершини x_j , якщо вона досягається з неї.

Матриця досяжності використана в контексті аналізу ієрархій для оцінки можливостей досягнення одних елементів через інші в системі зв'язків. Її формування пов'язане з теорією графів і аналізом зв'язності. Вона показує, чи можна перейти від одного елемента ієрархії до іншого (безпосередньо або через проміжні елементи).

У реалізованому нами варіанті прийнято спосіб безпосередніх зв'язків між факторами. Вказана матриця стала основою для подальшого виділення рівнів ієрархії та виявлення залежностей між елементами [57-59].

Квадратна бінарна матриця досяжності, що відповідає семантичній мережі,

має вигляд, відображений таблицею 2.1.

Таблиця 2.1

Матриця досяжності факторів достовірності текстових даних

	X1	X2	X3	X4	X5	X6	X7	X8
X1	1	1	1	0	1	0	1	1
X2	0	1	0	0	1	0	1	1
X3	0	0	1	0	1	1	1	1
X4	0	1	1	1	1	1	1	1
X5	0	0	0	0	1	0	1	0
X6	0	0	0	0	1	1	1	0
X7	0	0	0	0	0	0	1	1
X8	0	0	0	0	0	0	0	1

На основі матриці сформовано масив досяжних вершин $D(w_i)$, який кожному номеру рядка матриці i ставить у відповідність номери стовпців j , в яких розміщені одиниці. Вершин попередниці утворюють масив $P(w_i)$, котрий номеру стовпця j ставить у відповідність номери рядків i з одиницями.

У результаті перетин елементів масивів, тобто масив

$$Z(w_i) = D(w_i) \cap P(w_i), \quad (2.12)$$

ідентифікує певний рівень переваг факторів, що належать певним вершинам. Для реалізації (2.12) необхідним є виконання рівності

$$P(w_i) = Z(w_i). \quad (2.13)$$

На основі правил формування масивів даних $D(w_i)$ і $P(w_i)$ та умови (2.12) сформовано початкову таблицю ітераційного процесу визначення рівнів

пріоритетності факторів. У такому випадку співвідношення (2.13) свідчить про те, що номери факторів, зазначені у другому та третьому стовпцях таблиці, збігаються, що, у свою чергу, визначає сформований рівень їхньої вагомості.

Таблиця 2.2

Дані першого рівня пріоритетності факторів

i	$D(w_i)$	$P(w_i)$	$D(w_i) \cap P(w_i)$
1	1,2,3,5,7,8	1	1
2	2,5,7,8	1,2,3	2
3	3,5,6,7,8	1,3,4	3
4	2,3,4,5,6,7,8	4	4
5	5,7	1,2,3,4,5,6	5
6	5,6,7	3,4,6	6
7	7,8	1,2,3,4,5,6,7	7
8	8	1,2,3,4,7,8	8

Відповідно до умови (2.13) однакові номери зафіксовано у третьому і четвертому стовпцях табл. 2.2 для факторів 1 (професіоналізм автора) і чотири (джерело інформації), що було визначено згідно використаного методу моделювання ієрархій чинниками найвищого рівня. Далі вилучено з табл. 2.2 перший і четвертий рядки, а в третьому стовпці цифри 1 і 4. Одержано табл. 2.3.

Таблиця 2.3

Дані другого рівня пріоритетності факторів

i	$D(w_i)$	$P(w_i)$	$D(w_i) \cap P(w_i)$
2	2,5,7,8	2,3	2
3	3,5,6,7,8	3	3
5	5,7	2,3,5,6	5
6	5,6,7	3,6	6
7	7,8	2,3,5,6,7	7
8	8	2,3,7,8	8

Таблиця 2.3 забезпечує формування другого рівня для фактора 3 (контекст і структура тексту). Аналогічно попереднім діям приходимо до таблиці

наступного рівня пріоритетності факторів.

Таблиця 2.4

Дані третього рівня пріоритетності факторів

i	$D(w_i)$	$P(w_i)$	$D(w_i) \cap P(w_i)$
2	2,5,7,8	2	2
5	5,7	2,5,6	5
6	5,6,7	6	6
7	7,8	2,5,6,7	7
8	8	2,7,8	8

Третій рівень утворили фактори 2 (об'єктивність автора) і 6 (багаторазове публікування). Далі отримано табл. 2.5 четвертого рівня.

Таблиця 2.5

Дані четвертого рівня пріоритетності факторів

i	$D(w_i)$	$P(w_i)$	$D(w_i) \cap P(w_i)$
5	5,7	5	5
7	7,8	5,7	7
8	8	7,8	8

Четвертий рівень – фактор 5 (перевірка фактів). Без таблиць п'ятий рівень – фактор 7 (спростування і критика) і шостий – фактор 8 (соціальна довіра.)

У результаті запроектовано базову модель впливових факторів рис.2.2)

Базова модель рис. 2.2, отримана на основі методу аналізу ієрархій, вважається попереднім результатом, оскільки, по-перше, поки-що не отримано вагових показників переваг факторів, по друге – відсутнє врахування додаткових чинників, «захованих» в атомарних предикатах, вплив яких на існуючі зв'язки доцільно врахувати. Подальше дослідження орієнтоване на сукупне використання методу ранжування факторів і методу визначення вагомості предикатів для отримання модифікованої моделі пріоритетного впливу факторів на достовірність даних.



Рис. 2.2. Базова модель факторів достовірності інформаційних повідомлень

2.4. Модель факторів достовірності текстових повідомлень на основі методу ранжування та методу визначення вагомості предикатів

Передумовою отримання уточнених рівнів (рангів) є розрахунок числових вагових пріоритетів факторів без огляду на наявність предикатів. Для синтезування вказаної моделі наведемо коротко математичне трактування методу ранжування [60-62].

Нехай z_{ij} – кількість впливів чи залежностей для i -го типу зв'язку та j -го фактора ($j = 1, \dots, n$); w_i – вага i -го типу зв'язку. Прийmemo, що $i = 1$ визначає вплив фактора-джерела на фактор-приймач, $i = 2$ – залежність фактора від фактора. Крім того, для впливів ваги будуть додатними, тобто $w_1 > 0$, для залежностей – від'ємними, тобто $w_2 < 0$. Підсумкові вагові величини впливу факторів на достовірність тестових даних з урахуванням типів зв'язків позначено через змінну S_{ij} , значення якої отримано з виразу

$$S_{ij} = \sum_{i=1}^2 \sum_{j=1}^n z_{ij} w_i, \quad (2.14)$$

де n – умовний номер фактора [24].

Оскільки згідно з висловленими передумовами $S_{2j} < 0$, тому отримані при розрахунках значення змінено на величину $\Delta_j = \max|S_{2j}|$, ($j = 1, 2, \dots, n$), після чого формула (2.14) отримала такий вигляд:

$$S_{Fj} = \sum_{i=1}^2 \sum_{j=1}^n (z_{ij} w_i + \max|S_{2j}|), \quad (2.15)$$

де S_{Fj} – підсумкова величина числової міри впливу фактора на процес.

Кількісні показники для z_{ij} представлено у вигляді таблиці 2.6.

Таблиця 2.6

Назва фактора	Позначення	Кількість впливів	Кількість залежностей
Професіоналізм автора	X1	5	0
Об'єктивність автора	X2	3	2
Інформативність контенту	X3	3	2
Джерело інформації	X4	6	0
Перевірка фактів	X5	1	5
Багаторазове публікування	X6	2	2
Спростування і критика	X7	1	6
Соціальна довіра	X8	0	5

Прийmemo такі значення для вагових величин обох типів зв'язків: $w_1 = 10$., $w_2 = -10$ умовних одиниць. Результати розрахунку, виконані за формулою (2.15), занесено у таблицю 2.7.

Таблиця 2.7

Розрахункові дані попереднього ранжування факторів

j	z_{1j}	z_{2j}	S_{1j}	S_{2j}	S_{Fj}	Рівні
1	5	0	50	0	110	2
2	3	2	30	-20	70	3
3	3	2	30	-20	70	3
4	6	0	60	0	120	1
5	1	5	10	-50	20	5
6	2	2	20	-20	60	4
7	1	6	10	-60	10	6
8	0	5	0	-50	10	6

Звернемо увагу на особливість обчислення величини S_{Fj} в табл. 2.7. Величина $\Delta_j = \max|S_{2j}|=60$ обумовлює збільшення реальних значень змінної S_{Fj} на 60 умовних одиниць для отримання додатних величин при встановленні попередніх числових пріоритетів факторів.

Наступний крок стосується математичного трактування алгоритму визначення вагомості предикатів [60], що забезпечує встановлення остаточних рівнів та відповідних їм вагових значень пріоритетності факторів. Для цього введено показник впливу предикату у вигляді коефіцієнта вагомості k_{ip} , що визначає посилення чи послаблення зв'язків між факторами для i -го типу зв'язку та p -го предикату.

Нагадаємо, що $i = 1$ ідентифікує вплив одного фактора на інший, $i = 2$ – залежність, як тип зв'язку. Змінною l позначимо номер предикату.

Значення коефіцієнтів вагомості предикатів, обумовлені певними експертними пропозиціями, задано окремою таблицею.

Таблиця 2.8

Коефіцієнти вагомості предикатів семантичної мережі

l	Предикати впливу	k_{1,p_l}	Предикати	k_{2,p_l}
1	визначає	5	визначається	5
2	формує	5	формується	5
3	обумовлює	4	обумовлюється	4
4	стає основою	5	ґрунтується	5
5	передбачає	3	передбачається	3
6	враховує	3	враховується	3
7	діє	4	отримує	4

Крім інформації табл. 2.8, додатково використано такі вихідні дані: кількість і тип предикатів для кожного з факторів відображено у семантичній мережі рис. 2.1; лінгвістичну суть предикатів відтворено у формалізованому описі мережі. Таким чином, маємо достатньо повну вихідну базу даних для формування для кожного з факторів множини, елементи якої відповідають коефіцієнтам вагомості предикатів. Числові значення елементів цих множин забезпечили розрахунок вагової ідентифікації предикатів та підсумкових вагових пріоритетів факторів. Множини коефіцієнтів вагомості предикатів позначено через WIJ , де I визначає тип зв'язку, J – умовний номер фактора. Для типу зв'язку «вплив» узагальнений запис множин вагомості предикатів для факторів семантичної мережі виражено такими відношеннями:

$$\begin{aligned}
 x_1 \subset W11 &= \{k_{1,p_2}; k_{1,p_3}; k_{1,p_5}; k_{1,p_7}; k_{1,p_8}\}; x_2 \subset W12 = \{k_{1,p_5}; k_{1,p_7}; k_{1,p_8}\}; \\
 x_3 \subset W13 &= \{k_{1,p_5}; k_{1,p_6}; k_{1,p_7}; k_{1,p_8}\}; x_4 \subset W14 = \{k_{1,p_2}; k_{1,p_3}; k_{1,p_5}; k_{1,p_6}; k_{1,p_7}; k_{1,p_8}\}; \\
 x_5 \subset W15 &= \{k_{1,p_7}\}; x_6 \subset W16 = \{k_{1,p_5}; k_{1,p_7}\}; x_7 \subset W17 = \{k_{1,p_8}\}; x_8 \subset W18 = \{0\}
 \end{aligned}
 \quad (2.16)$$

Відповідно до значень коефіцієнтів вагомості з табл. 2.8, отримано:

$$\begin{aligned}
x_1 \in W11 &= \{4;5;3;3;5\}; x_2 \in M12 = \{5;4;4\}; x_3 \in W13 = \{5;3;5\}; \\
x_4 \in W14 &= \{4;4;3;5;5;5\}; x_5 \in W15 = \{4\}; x_6 \in W16 = \{3;5\}; \\
x_7 \in W17 &= \{5\}; x_8 \in W18 = \{0\}.
\end{aligned} \tag{2.17}$$

Для типу зв'язку «залежність» опустимо вирази подібно, як у (2.16) і визначимо множини (2.18), використовуючи опис семантичної мережі та числові ваги предикатів. При цьому для WIJ маємо: $I = 2$; $J = 1, \dots, 8$.

$$\begin{aligned}
x_1 \in W21 &= \{0\}; x_2 \in W22 = \{4;4\}; x_3 \in W23 = \{5;4\}; \\
x_4 \in W24 &= \{0\}; x_5 \in W25 = \{3;5;5;3;3\}; x_6 \in W26 = \{3;3\}; \\
x_7 \in W27 &= \{3;4;4;5;4;5\}; x_8 \in W28 = \{5;4;5;5;5\}.
\end{aligned} \tag{2.18}$$

Для зручності використання множин (2.17) і (2.18) для кожного з факторів визначено середні значення з урахуванням всіх компонент множин. Вони виражають узагальнені коефіцієнтами посилення чи послаблення дії фактора при попарній взаємодії між собою та забезпечать їх інтегральний вплив на достовірність текстових даних. Відповідні розрахунки виконано з урахуванням числових значень, заданих множинами (2.17) і (2.18).

Остаточні обчислення здійснено за формулою, логіка якої стосовно усереднення компонентів множин описана вище.

$$d_{ij} = \sum_{r=1}^{z_{ij}} (Wij_r / z_{ij}), \text{ для } i = 1, 2; j = 1, 2, \dots, 8, \tag{2.19}$$

де: Wij_r – елементи множин (2.17) і (2.18); z_{ij} – кількість елементів у множинах.

Вагові значення факторів P_{ij} отримано після множення вагових пріоритетів

S_{ij} таблиці 2.7 на коефіцієнти d_{ij} . Оскільки для типів зв'язку «залежність» $P_{2j} < 0$ поступаємо (як і вище) таким чином. Отримані числові результати зміщуємо на величину $\max |P_{2j}|$, ($j = 1, 2, \dots, 8$).

Остаточно отримано вираз для обчислення вагових пріоритетів факторів:

$$P_{Fj} = INT \left(\sum_{i=1}^2 \sum_{j=1}^8 (d_{ij} S_{ij} + \max |P_{2j}|) \right). \quad (2.20)$$

Вагові переваги факторів, розраховані за формулою (2.20) з використанням даних таблиці 2.7 та впливу предикатів, забезпечили чергову ітерацію процесу ранжування факторів, оформлену таблицею 2.9.

Таблиця 2.9

Ранжування факторів достовірності текстових даних
з урахуванням предикатів семантичних мереж

j	d_{1j}	d_{2j}	S_{1j}	P_{1j}	S_{2j}	P_{2j}	P_{Fj}	Ранг	Рівень пріор.
1	4.2	0	50	210	0	0	450	7	2
2	2.6	4	30	78	-20	-46	272	5	4
3	4.3	4.5	30	129	-20	-90	279	6	3
4	4.3	0	60	258	0	0	498	8	1
5	4	3.8	10	40	-50	-190	90	3	6
6	2	3	20	40	-20	-60	260	4	5
7	5	4.2	10	50	-60	-252	38	2	7
8	0	4.8	0	0	-50	-240	0	1	8

Підсумкову багаторівневу модель факторів впливу на достовірність текстових даних після етапу врахування дії предикатів семантичної мережі отримано з урахуванням колонки «Рівень пріоритетності» (рис.2.3).

З таблиці випливає, що високий (за цифрою) ранг фактора обумовлює відповідну пріоритетність, рівень якої встановлено, починаючи з одиниці.



Рис. 2.3. Багаторівнева модель факторів достовірності текстових повідомлень з врахуванням предикатів семантичної мережі

Модель отримана у результаті чергової ітерації процесу оцінювання дії множини чинників на показник достовірності текстових повідомлень з урахуванням впливу предикатів семантичної мережі. Попередні етапи визначення рівнів важливості факторів також реалізовані на базі семантичної мережі та використовували: метод аналізу ієрархій, в основі якого матриця досяжності; метод ранжування, побудований на підставі врахування кількостей типів зв'язків між факторами та їх умовних вагових значень. Перераховані засоби суттєво залежать від експертного оцінювання використаних параметрів, тому не можуть вважатися вирішальними щодо остаточних вагових пріоритетів факторів у процесі оцінювання показника достовірності текстових інформаційних повідомлень. З огляду на сказане здійснено оптимізацію вагових значень факторів з використанням методу попарних порівнянь.

2.5. Оптимізація моделі факторів формування показника достовірності текстових інформаційних повідомлень

Для оптимізації моделі факторів формування показника достовірності текстових повідомлень використано метод попарних порівнянь [61], який, як відзначено у [46], «встановлює кардинальну погодженість стосовно рівня пріоритетності факторів». Метод попарних порівнянь вважається важливим інструментом, що використовується для оцінювання переваг чинників досліджуваного процесу за допомогою їх порівнянь парами щодо певного критерію чи властивості. Цей метод особливо корисний у випадках, коли потрібно визначити передбачуваність чи перевагу одного фактора над іншими.

Основні кроки методу містять формування всіх можливих пар чинників, їх порівняння за заданим критерієм та надання оцінок або ранжування кожній парі відповідно до переваги вибору одного варіанту над іншим. Результати порівнянь агрегуються для отримання загального ранжування всіх альтернатив (чинників). Вказаний метод є досить універсальним у багатьох сферах, де потрібно об'єктивно оцінювати і порівнювати альтернативи. У результаті для кожного з факторів розраховуються числові пріоритети переваг, що відповідатимуть мірі їх впливу на ступінь достовірності даних. З іншого боку, отримані величини обумовлюють синтезування відповідної оптимізованої моделі переважаючої дії факторів на процес формування показника достовірності текстових повідомлень.

Відповідно до процедурної схеми обраного методу сформовано квадратну матрицю попарних порівнянь A , що є обернено-симетричною та містить числові значення, які відображають результати порівняння рівнів значущості факторів, представлених у структурі моделі (рис. 2.3). Значення елементів цієї матриці визначено на основі шкали відносної важливості (табл. 2.10) [61]. Зокрема, оцінку важливості між двома конкретними факторами розміщено у комірці на перетині відповідного рядка та стовпця, що позначено змінними, які представляють ці фактори. Водночас значення на головній діагоналі мають сталу величину, рівну

одиниці, що відповідає оцінці фактора.

Таблиця 2.10

Шкала відносної важливості об'єктів

Оцінка	Критерії порівняння	Пояснення щодо критерію
1	Об'єкти рівноцінні	Відсутність переваги x_1 над x_2
3	Один об'єкт дещо переважає інший	Слабка перевага x_1 над x_2
5	Один об'єкт переважає інший	Суттєва перевага x_1 над x_2
7	Об'єкт значно переважає інший	Явна перевага x_1 над x_2
9	Об'єкт абсолютно переважає інший	Абсолютна перевага x_1 над x_2
2,4,6,8	Компромісні проміжні значення	Допоміжні порівняльні оцінки

Остаточно матрицю попарних порівнянь помістимо у табл. 2.11.

Таблиця 2.11

Матриця попарних порівнянь факторів достовірності повідомлень

	X1	X2	X3	X4	X5	X6	X7	X8
X1	1	4	2	1/2	5	4	8	8
X2	1/4	1	1/2	1/4	3	2	5	7
X3	1/2	2	1	1/3	3	4	6	7
X4	2	4	3	1	6	5	8	9
X5	1/5	1/3	1/3	1/6	1	1/2	2	4
X6	1/4	1/2	1/4	1/5	2	1	4	5
X7	1/8	1/5	1/6	1/8	1/2	1/4	1	2
X8	1/8	1/7	1/7	1/9	1/4	1/5	1/2	1

Аналіз матриці попарних порівнянь виконано із використанням

програмного засобу «Імітаційне моделювання в системному аналізі методом бінарних порівнянь» [58], візуальне представлення інтерфейсу якого подано на рис. 2.4. У результаті обчислень сформовано головний власний вектор пріоритетів, нормалізовані значення компонент якого відображають розподіл вагових коефіцієнтів факторів, що впливають на рівень достовірності текстової інформації.

Введіть число критеріїв: 8

Введіть назви критеріїв

№	1	2	3	4	5	6	7	8
назва	1	2	3	4	5	6	7	8

Задання експертних оцінок переваг критеріїв

	1	2	3	4	5	6	7	8
1	1	4	2	1/2	5	4	8	8
2	1/4	1	1/2	1/4	3	2	5	7
3	1/2	2	1	1/3	3	4	6	7
4	2	4	3	1	6	5	8	9
5	1/5	1/3	1/3	1/6	1	1/2	2	4
6	1/4	1/2	1/4	1/5	2	1	4	5
7	1/8	1/5	1/6	1/8	1/2	1/4	1	2
8	1/8	1/7	1/7	1/9	1/4	1/5	1/2	1

Вивід проміжних результатів

	BI	E	En	En1	En2
1	0	2,908	0,243	2,047	8,408
2	0	1,265	0,105	0,886	8,370
3	0,58	1,897	0,158	1,328	8,360
4	0,9	3,884	0,325	2,751	8,459
5	1,12	0,590	0,049	0,406	8,222
6	1,24	0,840	0,070	0,592	8,415
7	1,32	0,326	0,027	0,226	8,276
8	1,41	0,227	0,019	0,163	8,599

Результати методу

$\lambda_{\max}$ 8,38908994856573

ІП 0,055584278366533

ВП 0,039421474018818

Рис. 2.4. Результат розрахунку вагових пріоритетів факторів ДТП

Пояснимо коротко алгоритм отримання числових пріоритетів факторів, користуючись даними вікна інтерфейсу «Вивід проміжних результатів».

Компоненти головного власного вектора $A(a_1, a_2, \dots, a_n)$ розраховано через середні геометричні значення елементів рядків матриці за формулою

$$a_i = \sqrt[n]{a_{i1} \cdot a_{i2} \cdot \dots \cdot a_{in}} \quad i = \overline{1, n}. \quad (2.20)$$

В результаті отримано вектор

$$A = (2,908; 1,265; 1,897; 3,884; 0,590; 0,840; 0,326; 0,227),$$

нормалізація якого за формулою

$$a_{i \text{ норм}} = \frac{\sqrt[n]{a_{i1} \cdot a_{i2} \cdot \dots \cdot a_{in}}}{\sum_{i=1}^n \sqrt[n]{a_{i1} \cdot a_{i2} \cdot \dots \cdot a_{in}}} \quad i = \overline{1, n} \quad (2.21)$$

розрахунок головного власного вектора матриці, компоненти якого відтворені в опції E_n :

$$A_{\text{норм}} = (0,243; 0,105; 0,158; 0,325; 0,049; 0,070; 0,027; 0,019).$$

Як було відзначено вище, компоненти отримано вектора відображають прикінцевий результат. З іншого боку, більш зручним в користуванні буде вектор $A_{\text{норм}}$, компоненти якого помножені на коефіцієнт $k = 500$. Остаточного одержано такий масштабований остаточний вектор

$$A_{\text{норм}} \times k = (243; 105; 158; 325; 49; 70; 27; 19). \quad (2.22)$$

Нормальною практикою досліджень вважається перевірка узгодженості отриманих результатів за певними еталонами. Такими критеріями для даного методу є наступні: значення $\lambda_{\max}$ обернено-симетричної матриці; індекс узгодженості IU ; відношення узгодженості WU (опції $\lambda_{\max}$, ПІ, ВП вікна інтерфейсу). Для їх отримання використовуємо опцію E_{n1} : $A_{\text{норм}1}$

$$A_{\text{норм}1} = (2,047; 0,886; 1,328; 2,751; 0,406; 0,592; 0,226; 0,163).$$

Продовжуючи, компоненти вектора $A_{\text{норм}1}$ поділено на відповідні компоненти вектора $A_{\text{норм}}$, що обумовлює вектор $A_{\text{норм}2}$ (опція E_{n2}):

$$A_{\text{норм}2} = (8,408; 8,370; 8,360; 8,459; 8,222; 8,415; 8,276; 8,599).$$

Значення компонент вектора $A_{\text{норм}2}$ обумовлюють величину $\lambda_{\max}$. Згідно з програмою $\lambda_{\max} = 8,39$. Причому, ціла частина його значення не повинна перевищувати кількості факторів досліджуваного процесу. Більш точну оцінку

отримано внаслідок розрахунку величини узгодженості

$$IU = (\lambda_{\max} - n) / (n - 1). \quad (2.23)$$

Розраховане значення $IU = 0,056$ (величина ІІ у вікні «Результати методу»). Показник індексу узгодженості демонструє відповідність розрахованої величини нормативній величині, названій випадковим індексом, сформованим шляхом комп'ютерної генерації для матриць різних розмірів. Еталонний показник наведено у вигляді спеціалізованої табличної залежності [46].

Таблиця 2.12

Варіанти значень випадкового індексу

Кількість об'єктів	3	4	5	6	7	8	9	10	11	12
Еталонне значення	0,58	0,90	1,12	1,24	1,32	1,41	1,45	1,49	1,51	1,54

Результати отриманих розрахунків задовільні, якщо виконується умова $IU < 0,1 \times WI$. Згідно таблиці для восьми об'єктів $WI = 1,41$, що обумовлює $0,056 < 0,1 \times 1,41$. На завершення оцінюють величину відношення узгодженості, що отримується за виразом $WU = IU / WI$. Його значення оптимальне, якщо $WU \leq 0,1$ (величина ВІІ у вікні «Результати методу»). Розраховане $WU = 0,04$, що остаточно свідчить про достовірність результатів.

Огляд числових результатів, отриманих у процесі обробки матриці попарних порівнянь, засвідчив відповідність ключовим аналітичним принципам дослідження, зокрема щодо обґрунтованості отриманих вагових коефіцієнтів факторів, прийнятної точності обчислювальних процедур та внутрішньої логічної узгодженості експертних оцінок при здійсненні попарних порівнянь. Оптимальні вагові характеристики факторів, отримані на основі виразу (2.22), підтверджують дотримання принципу відповідності кожного фактора окремому рівню впливу, що забезпечує структурну однозначність у моделюванні. На підставі цього можна зробити додатковий висновок: за умови, що домінантне

власне значення матриці та відповідний індекс узгодженості не перевищують встановлених нормативних меж, ці параметри можуть бути розцінені як формальні підстави для побудови оптимізованої багаторівневої моделі визначення пріоритетного впливу факторів на показник достовірності текстових повідомлень, концептуальну структуру якої відображено на рис. 2.5.



Рис. 2.5. Остаточна модель впливу факторів на рівень достовірності текстових повідомлень

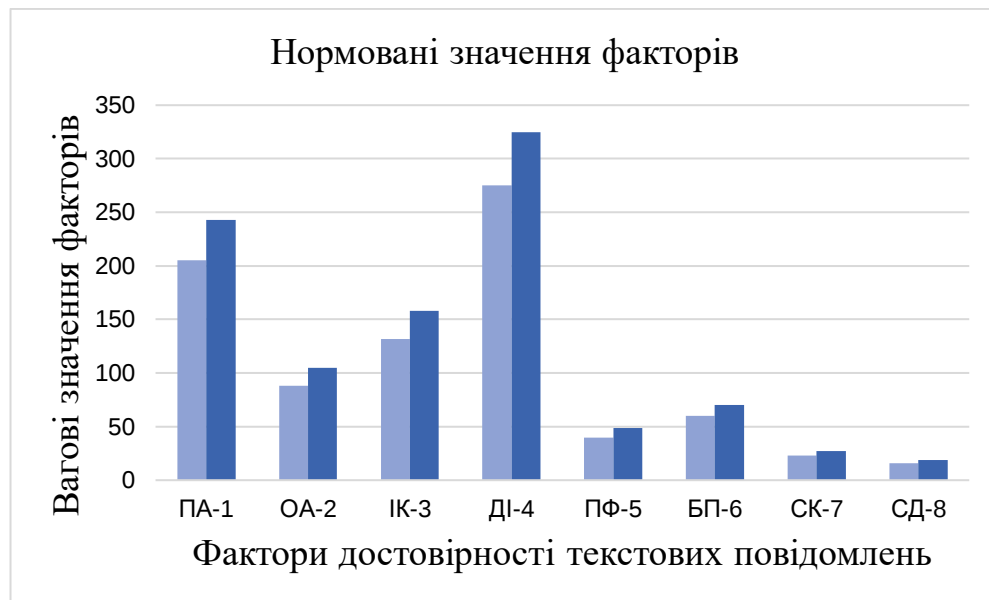
Синтезована модель рис. 2.5 теоретично підтверджує інтуїтивне розуміння переваги чинників «Джерело інформації» та «Професіоналізм автора» у процесі формування показника достовірності текстових повідомлень.

Результати застосування методу попарних порівнянь уможливили отримання оптимізованої моделі пріоритетного впливу факторів на показник достовірності текстових повідомлень, що обумовлює покращення об'єктивності і ефективності прийняття рішень у різних сферах життя і діяльності при дотриманні рекомендацій дослідження щодо аналізу отриманої інформації.

Порівняння умовних ваг пріоритетності факторів, отриманих на підставі методу ранжування (вектор Y_1) і методу попарних порівнянь (вектор $Y_2 = Y_{\text{норм}} \times k$ масштабований коефіцієнтом $k = 1000$), наведено у таблиці 2.13 та візуалізовано відображено на рис. 2.6.

Таблиця 2.13

Фактори	ПА-1	ОА-2	ІК-3	ДІ-4	ПФ-5	БП-6	СК-7	СД-8
Y_1	205	88	132	275	40	60	23	16
$Y_2 = Y_{\text{норм}} \times k$	243	105	158	325	49	70	27	19

Рис. 2.6. Візуалізація умовних ваг факторів за векторами (Y_1) та (Y_2)

Візуальне відтворення векторів Y_1 та Y_2 на підставі даних табл. 2.13 вказує на те, що компоненти вказаних векторів, отримані за різними методами, практично співпадають за вагами факторів у відповідних багаторівневих моделях. Сказане свідчить про можливість використання оптимізованої моделі рис 2.6 та компонент вектора Y_2 для проведення наступних досліджень, які полягатимуть у знаходженні альтернативних та оптимального варіантів реалізації процесу оцінювання достовірності текстових повідомлень.

2.6. Альтернативні та оптимальний варіанти процесу оцінювання рівня достовірності текстових повідомлень

Продовження дослідження дисертаційної теми в озвученому вище напрямі є актуальним завданням, оскільки з плином часу збільшуються обсяги та складність інформаційних повідомлень, що обумовлює важливість процесів, які забезпечують умови досягнення їх достовірності. Дослідження альтернативних та оптимальних методів оцінювання достовірності даних може сприяти покращенню якості аналізу, надійності та адекватності прийняття рішень на основі текстових даних, підвищенню рівня довіри до інформації, що циркулює у суспільстві. Водночас, вони є важливим кроком у розвитку методів обробки та аналізу інформації в цифровому світі.

Результати попередніх досліджень, присвячених плануванню та реалізації альтернатив у виконанні окремих процедур видавничо-поліграфічних процесів, а також їх узагальнення [63–65], дозволили сформулювати концептуальні засади для розроблення та удосконалення підходів до визначення достовірності текстових інформаційних повідомлень. Такий підхід здатен забезпечити зростання довіри користувачів до змісту отриманої інформації завдяки підвищенню якості її попередньої оцінки.

Сформовані в попередніх параграфах багаторівневі моделі, що описують ієрархію впливу окремих факторів на процес встановлення достовірності текстової інформації, створюють підґрунтя для обґрунтованого визначення значущості кожного з чинників та їхнього внеску у загальний рівень якості інформаційного повідомлення. З позицій математичного аналізу наявні вихідні параметри, що дозволяють оцінити ступінь впливу факторів, однак цієї інформації недостатньо для повноцінного управління процесом формування якісного показника достовірності. Вирішальним є не лише знання про умовну пріоритетність окремих чинників, а й розуміння витрат ресурсів, зокрема трудомісткості, необхідної для реалізації впливу кожного з них у межах

комплексного оцінювання з метою досягнення прогнозованого рівня достовірності текстових даних.

Описане завдання стосується процесу проєктування та аналізу можливих альтернатив з подальшим вибором оптимального варіанта відповідно до обраного критерію. Такі задачі належать до класу багатокритеріальної (у даному випадку — багатфакторної) оптимізації, що передбачає прийняття виважених рішень на основі множини параметрів. При цьому розв’язання подібних задач не потребує урахування повного переліку факторів — достатньо обмеженого набору чинників, релевантність яких визначається відповідно до принципу Парето [66,67].

Згідно з цим підходом, з загальної множини обираються ті фактори, вплив яких є переважаючим, а чинники з незначною вагомістю виключаються з подальшого аналізу. Як результат, дослідження базується на множині взаємно непомітованих факторів, що формують так звану множину Парето.

Відповідно до підходів, сформованих у межах теорії прийняття рішень [46,49], задача багатокритеріальної оптимізації на множині можливих альтернатив за наявності функцій мети зводиться до знаходження таких значень функцій корисності, які забезпечують їх максимізацію. У цьому контексті механізм вибору оптимального варіанта реалізується через процедуру лінійного агрегування критеріїв [68], що передбачає формування інтегрального показника шляхом поєднання окремих цільових функціоналів у вигляді їх лінійної комбінації.

$$F(w, x) = \sum_{i=1}^m w_i f_i(x) \rightarrow \max_{x \in D}; \quad w \in W, \quad (2.24)$$

$$\text{де } W = \left\{ w = (w_1, \dots, w_m)^T; \quad w_i > 0; \quad \sum_{i=1}^m w_i = 1 \right\}.$$

Кількісні значення вагових коефіцієнтів факторів інтерпретовано як числові вирази відповідних функцій корисності. Для здійснення вибору серед альтернативних варіантів доцільно застосувати положення теореми

багатокритеріальної теорії корисності [69], яка формулює, що за умови незалежності критеріїв як за корисністю, так і за перевагою, можна побудувати інтегральну функцію корисності, яка узагальнює вплив усіх критеріїв на процес прийняття рішення.

$$U(x) = \sum_{i=1}^m w_i u_i(y_i), \quad (2.25)$$

яка вважається показником оптимальності альтернативи. Загалом: $U(x)$ – функція корисності ($0 \leq U(x) \leq 1$) варіанту X ; $u_i(y_i)$ – величина корисності i -го критерію ($0 \leq u_i(y_i) \leq 1$); y_i – отримане значення варіанту X за критерієм i ; w_i – вага i -го критерію, причому $0 < w_i < 1$, $\sum_{i=1}^m w_i = 1$.

Узагальнюючи наведене, відначимо, що завдання вибору альтернативи розв’язано з використанням таких тверджень [70,71]:

- множина альтернатив X – складається з множини $X = \{x_1, x_2, \dots, x_n\}$, які можна досягнути;
- критеріями оцінювання альтернативних варіантів вважаються функції корисності $f_i : X \rightarrow R(i = \overline{1, m})$.
- приймаючий рішення суб’єкт використовує для прийняття остаточного варіанту ваговими значеннями факторів.

Спираючись на попередньо обґрунтовані теоретичні засади, для досліджуваного процесу виокремлюється множина Парето — підмножина факторів, жоден з яких не домінує над іншими, однак сукупно вони чинять найвагоміший вплив на процедуру оцінювання достовірності текстової інформації. Очевидно, що саме цим чинникам притаманні найвищі вагові коефіцієнти у відповідній моделі. Припустимо, що до цієї множини належать такі фактори Z_1, Z_2, Z_3, Z_4 .

Розглянемо три можливі варіанти реалізації процедури оцінювання текстових даних, які умовно позначимо як А1, А2, А3. Для кожного з них

визначено відповідні комбінації значень ступенів важливості з урахуванням впливу попередньо відібраних факторів. Хоча теоретично кількість таких комбінацій може бути досить великою, їхні фактичні значення залежать від особливостей конкретного інформаційного середовища.

Враховуючи формулу (2.25), отримаємо вхідні дані для проведення обчислень: $m = 4$; $u_i(y_i) = u_{ij}$ – корисність j -ї альтернативи ($j = 1, 2, 3$) за i -м фактором ($i = 1, \dots, 4$). У підсумку розрахунок значень функцій корисності буде реалізовано за такою формулою:

$$U_j = \sum_{i=1}^4 w_i u_{ij}; \quad j = 1, 2, 3, \quad (2.26)$$

де U_j – критерій корисності j -го альтернативного варіанту.

Результати оцінювання альтернативних варіантів на основі показників ефективності відповідних факторів подано у таблиці.

Таблиця 2.13

Вихідні дані для оцінювання альтернатив

Фактори	Ваги факторів	Оцінювання альтернатив за мірами ефективності факторів		
		A1	A2	A3
Фактор Z_1	w_{s1}	p_{11}	p_{12}	p_{13}
Фактор Z_2	w_{s2}	p_{21}	p_{22}	p_{23}
Фактор Z_3	w_{s3}	p_{31}	p_{32}	p_{33}
Фактор Z_4	w_{s4}	p_{41}	p_{42}	p_{43}

У табл. 2.13: w_{si} – стартова вага i -го фактора; p_{ij} – величина оцінювання факторів, відносно j -ї альтернативи.

Фактори, наведені в таблиці 2.13, утворюють початкову множину з відповідними ваговими коефіцієнтами, що були визначені на попередньому етапі. Спираючись на дані другого стовпця цієї таблиці та використовуючи шкалу Сааті (табл. 2.10), формуємо матрицю попарних порівнянь, яку для

зручності подано у вигляді таблиці (табл. 2.14). Оброблення цієї матриці дозволяє отримати уточнені вагові значення факторів, що входять до множини Парето, які надалі використано при розрахунку функцій корисності для альтернатив, що аналізуються.

Таблиця 2.14

Матриця попарних порівнянь факторів

	Z_1	Z_2	Z_3	Z_4
Z_1	1	w_{s1}/w_{s2}	w_{s1}/w_{s3}	w_{s1}/w_{s4}
Z_2	w_{s2}/w_{s1}	1	w_{s2}/w_{s3}	w_{s2}/w_{s4}
Z_3	w_{s3}/w_{s1}	w_{s3}/w_{s2}	1	w_{s3}/w_{s4}
Z_4	w_{s4}/w_{s1}	w_{s4}/w_{s2}	w_{s4}/w_{s3}	1

Позначення w_{si}/w_{sj} у матриці означають порівняння значень w_{si} та w_{sj} .

Після опрацювання матриці розраховано повторно вагові значення w_1, w_2, w_3, w_4 факторів – нової Парето-множини. Для контролювання результату нормалізації служать такі критерії: максимальне власне значення $\lambda_{\max}$; індекс узгодженості IU ; відношення узгодженості WU . Параметри значень вказаних критеріїв наведено у попередньому параграфі.

Для отримання значень оцінок корисності кожної альтернативи за i -м фактором будуюмо матриці попарних порівнянь переваг альтернатив стосовно факторів. Так, для фактора Z_1 отримаємо таку матрицю-таблицю.

Z_1	A1	A2	A3
A1	1	p_{11}/p_{12}	p_{11}/p_{13}
A2	p_{12}/p_{11}	1	p_{12}/p_{13}
A3	p_{13}/p_{11}	p_{13}/p_{12}	1

Елементи матриці вказують на результати порівняння мір ефективності фактора Z_1 залежно від варіантів.

За результатами опрацювання матриці отримуємо функції корисності

альтернатив для $Z_1: u_{11}, u_{12}, u_{13}$.

Аналогічно отримуємо функції корисності для факторів: $Z_2 - u_{21}, u_{22}, u_{23}$;
 $Z_3 - u_{31}, u_{32}, u_{33}$; $Z_4 - u_{41}, u_{42}, u_{43}$..

Узагальнені багатofакторні значення функцій корисності для альтернатив A1, A2, A3 обчислюються відповідно до формули (2.26) із врахуванням таких співвідношень:

$$\begin{aligned} U_1 &= w_1 \cdot u_{11} + w_2 \cdot u_{21} + w_3 \cdot u_{31} + w_4 \cdot u_{41}; \\ U_2 &= w_1 \cdot u_{12} + w_2 \cdot u_{22} + w_3 \cdot u_{32} + w_4 \cdot u_{42}; \\ U_3 &= w_1 \cdot u_{13} + w_2 \cdot u_{23} + w_3 \cdot u_{33} + w_4 \cdot u_{43}. \end{aligned} \quad (2.27)$$

Оптимальним серед запропонованих альтернатив прийнято вважати варіант, для якого об'єднана функція корисності, що інтегрує часткові цільові функціонали, набуває максимального значення. Згідно з [72] вона є критерієм прийняття рішення щодо доцільності реалізації відповідного варіанту в процесі оцінювання достовірності текстових повідомлень. Для реалізації цього підходу використано методологію багатofакторного аналізу з урахуванням процедур аналізу ієрархій і попарного порівняння, що сприяє зниженню рівня суб'єктивності оцінок та забезпечує вищу точність результатів.

Практичну реалізацію теоретичних викладок виконано для чотирьох факторів оптимізованої моделі рис. 2.5. Вказані чинники утворюють нову множину Парето, тому згідно вище наведених тверджень змінимо їх формалізовані позначення. Отже, для подальшого дослідження маємо фактори:

- $X_4(R_1)$ – джерело інформації (вага фактора 325 у. о.);
- $X_1(R_2)$ – професіоналізм автора (вага фактора 243 у. о.);
- $X_3(R_3)$ – інформативність контенту (вага фактора 158 у. о.);
- $X_2(R_4)$ – об'єктивність автора (вага фактора 105 у. о.).

Оскільки ми вибрали три альтернативні варіанти, задаємо комбінації величин ефективності факторів [73].

Таблиця 2.15

Комбінації значень ефективності факторів в альтернативах

Комбінації величин ефективності фактора у відсотках			
10 - 10 - 80	20 - 10 - 70	30 - 10 - 60	40 - 10 - 50
10 - 20 - 70	20 - 20 - 60	30 - 20 - 50	40 - 20 - 40
10 - 30 - 60	20 - 30 - 50	30 - 30 - 40	40 - 30 - 30
10 - 40 - 50	20 - 40 - 40	30 - 40 - 30	40 - 40 - 20
10 - 50 - 40	20 - 50 - 30	30 - 50 - 20	40 - 50 - 10
10 - 60 - 30	20 - 60 - 20	30 - 60 - 10	50 - 10 - 40
10 - 70 - 20	20 - 70 - 10	60 - 10 - 30	50 - 20 - 30
10 - 80 - 10	70 - 10 - 20	60 - 20 - 20	50 - 30 - 20
80 - 10 - 10	70 - 20 - 10	60 - 30 - 10	50 - 40 - 10

Почнемо з таблиці 2.16 .

Таблиця 2.16

Оцінювання альтернатив за вибраними факторами

Назви факторів	Ваги факторів (початкові)	Оцінювання альтернатив за мірами ефективності факторів		
		A1	A2	A3
Джерело інформації (R_1)	325 (w_{s3})	70	10	20
Професіоналізм автора (R_2)	243 (w_{s1})	20	50	30
Інформативність контенту (R_3)	158 (w_{s4})	40	20	40
Об'єктивність автора (R_4)	105 (w_{s2})	10	30	60

Для розрахунку функцій корисності факторів табл. 2.16 побудовано матрицю попарних порівнянь (табл. 2.17). Спочатку визначаємо ступінь значущості кожного фактора відносно інших, використовуючи метод попарних порівнянь. Потім обчислюємо відношення пріоритетів, що дозволяє отримати нормалізовані вагові коефіцієнти.

Таблиця 2.17

Матриця попарних порівнянь факторів альтернатив

	R_1	R_2	R_3	R_4
R_1	1	3	6	7
R_2	1/3	1	4	6
R_3	1/6	1/4	1	4
R_4	1/7	1/6	1/4	1

Опрацювання матриці за програмою [19] зумовило такі нормалізовані вагові значення факторів, задіяних у проектуванні альтернативних варіантів:

$$w_1 = 0,56; w_2 = 0,28; w_3 = 0,11; w_4 = 0,05.$$

Процес опрацювання матриці завершився при таких значеннях критеріїв:
 $\lambda_{\max} = 4,24; IU = 0,08; WU = 0,09.$

Отримані дані вважаємо достовірними, оскільки значення $\lambda_{\max}$, індекс узгодженості IU та відношення узгодженості WU знаходяться в межах норми.

Виконаємо розрахунок функцій корисності u_{ij} . Для цього будуємо матриці попарних порівнянь. Для фактора R_1 (джерело інформації) формуємо таку матрицю:

R_1	A1	A2	A3
A1	1	7	4
A2	1/7	1	1/3
A3	1/4	3	1

Отримано: $\lambda_{\max} = 3,03; IU = 0,02; WU = 0,03$. Корисність варіантів стосовно чинника R_1 : $u_{11} = 0,70; u_{12} = 0,08; u_{13} = 0,21$.

Для інших факторів наведемо лише кінцеві результати.

Фактор R_2 – професіоналізм автора:

$$\lambda_{\max} = 3,05; IU = 0,03; WU = 0,05; u_{21} = 0,13; u_{22} = 0,66; u_{23} = 0,21.$$

Фактор R_3 – інформативність контенту:

$$\lambda_{\max} = 3,00; IU = 0,00; WU = 0,00; u_{31} = 0,44; u_{32} = 0,11; u_{33} = 0,44.$$

Фактор R_4 – об'єктивність автора:

$$\lambda_{\max} = 3,05; IU = 0,03; WU = 0,05; u_{41} = 0,08; u_{42} = 0,27; u_{43} = 0,64.$$

Отримані дані, як видно з розрахунків, відповідають прийнятим критеріям. Остаточно числові значення ваг факторів та функцій корисності, отримані вище, обумовили таке остаточне числове трактування виразів (2.27)

$$\begin{aligned} U_1 &= 0.56 \cdot 0.70 + 0.28 \cdot 0.13 + 0.11 \cdot 0.44 + 0.05 \cdot 0.08; \\ U_2 &= 0.56 \cdot 0.08 + 0.28 \cdot 0.66 + 0.11 \cdot 0.11 + 0.05 \cdot 0.27; \\ U_3 &= 0.56 \cdot 0.21 + 0.28 \cdot 0.21 + 0.11 \cdot 0.44 + 0.05 \cdot 0.64. \end{aligned} \quad (2.28)$$

Після обчислень отримано: $U_1 = 0.480$; $U_2 = 0.256$; $U_3 = 0.257$.

Підсумкові числові характеристики часткових цільових функціоналів визначають пріоритетність альтернативних варіантів процесу формування показника достовірності текстових повідомлень.

Оптимальним вважаємо варіант A1, функціонал якого $U_1 = 0.480$ отримав максимальне значення, який прогнозовано при заданих частках ефективності факторів забезпечує належний рівень достовірності текстових даних і, як наслідок, формує належну довіру до отриманої інформації. Його перевагу забезпечила визначальна «участь» фактора «Джерело інформації».

Великі обсяги даних, швидкість їх розповсюдження та ризики дезінформації зумовлюють необхідність розроблення ефективних методів автоматизованого аналізу достовірності та врахування факторів множини Парето, зокрема джерела інформації, семантичного змісту контенту та структурних характеристик тексту, професіоналізму та об'єктивності автора.

Додатковим результатом досліджень даного розділу є розроблена узагальнена функціональна модель алгоритму автоматизованого визначення оптимального варіанту процесу оцінювання достовірності текстових

інформаційних повідомлень. Розроблена модель забезпечує можливість формалізованого представлення аналізу Парето-критеріїв та їхнього впливу на процес оцінювання достовірності повідомлень, передбачаючи як формування множини альтернатив, так і обчислення оптимального варіанту шляхом застосування методу лінійної агрегації критеріїв (рис. 2.7)

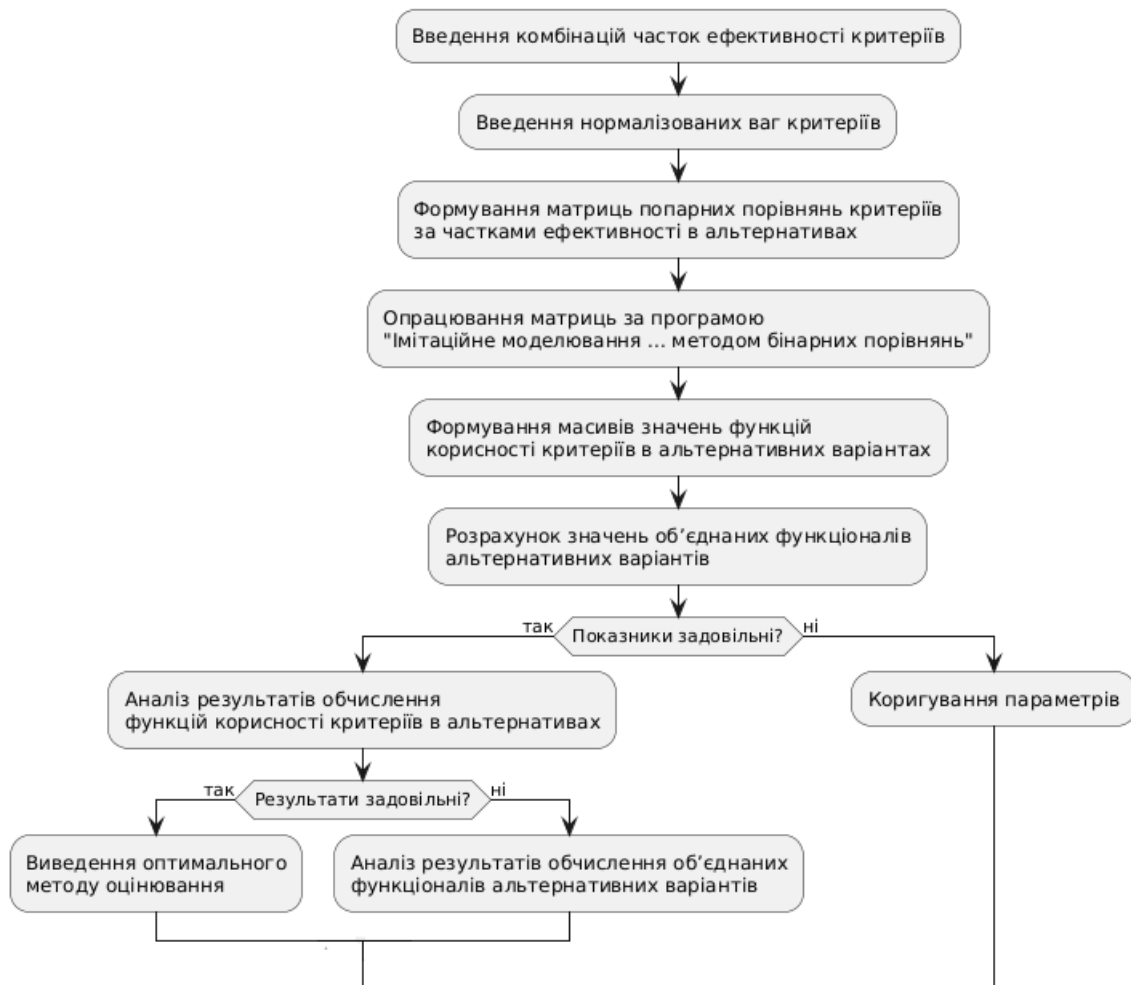


Рис. 2.7. Блок-схема визначення оптимального варіанту процесу оцінювання достовірності текстових інформаційних повідомлень

На завершення звернемо увагу на особливості застосування запропонованого підходу. Механізм формування вхідних даних, зокрема факторів, що впливають на правдивість контенту повідомлень, дає змогу здійснювати оцінювання рівня достовірності повідомлень на основі аналізу чинників впливу на процес моніторингу медійного простору.

Висновки до розділу 2

1. Проведено аналіз та описано чинники, що стосуються визначення ступеня вірогідності текстових інформаційних повідомлень про процес, або в більш загальному розумінні рівня достовірності текстових повідомлень, що уможливило їх виокремлення за типами та мірою впливу на правдивість контенту повідомлень і групування у вигляді узагальнених факторів.

2. Розроблено графічно орієнтовану модель у вигляді семантичної мережі, що уможливила відтворення формалізованого відображення структурних та лінгвістичних відношень між факторами достовірності текстових інформаційних повідомлень. Виконано опис семантичної мережі з використанням елементів логіки предикатів, що забезпечує можливість застосування математичного апарату теорії моделювання для вивчення особливостей впливу вказаних чинників на процес формування показника достовірності текстових повідомлень.

3. Запроектовано базову багаторівневу модель впливових факторів достовірності текстових інформаційних повідомлень за методом математичного моделювання ієрархій, в основі якої ітераційний процес покрокового формування рівнів пріоритетності факторів, реалізований на підставі побудови та опрацювання матриці досяжності, що формалізує наявність залежностей між факторами у семантичній мережі.

4. Розроблено на основі сукупного використання методів ранжування та визначення вагомості предикатів удосконалену багаторівневу модель факторів достовірності текстових повідомлень, поетапний процес створення якої передбачає: отримання варіанту на основі числа впливів (з вагою $w_1 = 10$) та залежностей (з вагою $w_2 = -10$) між факторами у семантичній мережі; коригування результату через врахування умовних вагових коефіцієнтів предикатів k_{ip} , які ідентифікують посилення чи послаблення залежностей між факторами.

5. Виконано оптимізацію моделі, отриманої на попередньому етапі. Опрацювання запроектованої матриці обумовило отримання покращених вагових переваг факторів за відповідними критеріями: $\lambda_{\max} = 8,39$; $IU = 0,056$; $WU = 0,04$. За отриманими даними синтезовано оптимізовану графічну модель переважаючої дії факторів на хід формування показника достовірності текстових інформаційних повідомлень.

6. Здійснено математичну постановку процесу дослідження альтернативних та оптимальних методів оцінювання рівня достовірності даних, що сприятиме покращенню якості аналізу, надійності та адекватності прийняття рішень на основі текстових даних, підвищенню рівня довіри до інформації, що циркулює у суспільстві.

7. Запроектовано на підставі методу лінійного згортання критеріїв альтернативні варіанти визначення рівня довіри до інформації з використанням функцій корисності, часток ефективності факторів множини Парето в альтернативах та часткових функціоналів. Розраховано та визначено серед альтернативних оптимальний варіант процесу формування показника достовірності текстових повідомлень за критерієм максимального значення цільового функціоналу, величина якого $U_1 = 0.480$.

8. Розроблено узагальнену функціональну модель алгоритму автоматизованого визначення оптимального варіанту процесу оцінювання достовірності текстових інформаційних повідомлень. Запропонована модель обумовлює формалізацію процесу аналізу Парето-критеріїв з огляду на генерування альтернативних та розрахунків оптимального варіантів процесу формування показника достовірності текстових інформаційних повідомлень.

РОЗДІЛ 3. МЕТОДИ ФОРМУВАННЯ ПОКАЗНИКА ОЦІНЮВАННЯ ДОСТОВІРНОСТІ ТЕКСТОВИХ ІНФОРМАЦІЙНИХ ПОВІДОМЛЕНЬ

3.1. Основні поняття і засоби теорії нечіткої логіки

Попередній розділ завершив дослідження основних чинників впливу на достовірність текстових повідомлень розробленням моделей, що відтворюють загальний алгоритм реалізації початкових етапів інформаційного підходу до вирішення вказаної проблематики. Остаточною завданням полягає в отриманні числового значення показника прогностичного оцінювання рівня правдивості та обґрунтованості отриманих повідомлень. Особливістю реалізації такого завдання полягає у тому, що в умовах сучасного суспільства обсяг інформації зростає експоненційно, а кількість неправдивих або маніпулятивних повідомлень та способів їх отримання також збільшуються. Підтвердженням цьому служить виконаний аналіз літературних джерел з даної тематики, наведений у першому розділі, який засвідчив зростаючий інтерес наукової спільноти, працівників сфери інформаційних послуг та великої аудиторії користувачів щодо проблеми встановлення та забезпечення надійності та правдивості контенту повідомлень.

Сказане обумовлює фіксування загального стану проблематики та наявності певного обсягу вихідних даних на момент дослідження. З іншого боку, зазначимо особливості досліджуваної тематики, що характеризується словом «невизначеність», яку традиційні алгоритми не завжди можуть врахувати в повній мірі. До них віднесемо нижче наведені узагальнені ознаки використаної нами вихідної інформаційної бази [74].

Текстові повідомлення можуть бути суб'єктивними, маніпулятивними чи викривленими.

Достовірність інформації може залежати від контенту або стилю викладення.

Текстова інформація часто містить неявні смисли, метафори або змістовні залежності.

Викликає сумнів кваліфікація та професіоналізм автора повідомлень.

Велике значення має освіта автора та знання відповідної галузі.

Наявність соціальної довіри щодо авторитетності джерела інформації. Відоме та надійне джерело збільшує достовірність інформації.

Наявність спростування і критики. Більшу достовірність має інформація, яка витримує критику і підтвердження даних.

Для вирішення вказаних і подібних проблем сучасна наука і практика пропонує використовувати засоби теорії нечітких множин, визначальним елементом якої є нечітка логіка як перспективний інструмент для аналізу текстових даних, оскільки дозволяє працювати з характерними для текстів невизначеністю, неповнотою та неоднозначністю інформації [75,76]. Нечітка логіка забезпечує роботу з даними, які важко формалізувати у вигляді строгих математичних моделей. Наприклад, поняття «достовірність» можна розглядати як змінну з градаціями (від «абсолютно достовірно» до «абсолютно недостовірно»). Системи, орієнтовані на нечітку логіку, здатні імітувати способи, на основі яких приймають рішення в умовах невизначеності. Це особливо важливо для аналізу текстових інформаційних повідомлень, де значна частина інформації є суб'єктивною. Водночас, нечітку логіку можна комбінувати з методами машинного навчання, обробки природної мови (NLP) або статистичного аналізу для досягнення кращих результатів. Більш детально про вказані варіанти сказано у першому розділі.

Найважливішим інгредієнтом нечіткої логіки прийнято вважати лінгвістичні змінні (ЛЗ), значення яких виражені словами або фразами (наприклад, «низький», «високий», «достовірний»). Кожне значення лінгвістичної змінної пов'язане з нечіткою множиною. Нечіткі множини визначаються функцією належності $\mu(x)$, яка приймає значення в інтервалі $[0,1]$. Визначальною особливістю лінгвістичних змінних є функції належності, що визначають, наскільки елемент відповідає заданій характеристиці. Прикладами функцій належності можуть бути трикутна, трапецієподібна, гаусівська. Функції належності сформовано на основі терм-множини значень та відповідних

лінгвістичних термінів факторів.

У роботах основоположника нечіткої логіки Заде [75] запроваджено стосовно деякої проблемної області поняття універсальної множини D . Зокрема нечітку підмножину M множини D сформовано за допомогою універсальної множини або відповідної шкали D і функції належності $\mu_M(d)$, тобто

$$M = \{(\mu_M(d), d), d \in D\}, \quad (3.1)$$

де $(0 \leq \mu_M(d) \leq 1)$.

Функція належності (ФН) визначає ступінь належності кожного елемента нечіткої множини до універсальної множини: $M \in D$.

За умови дискретності і скінченності базової шкали (тобто поділеної на кванти або частини чи проміжки) нечітка множина M може бути подана як

$$M = (\mu_M(d_1) / d_1, \mu_M(d_2) / d_2, \dots, \mu_M(d_n) / d_n) = \sum_{i=1}^n \mu_M(d_i) / d_i, \quad (3.2)$$

або у спрощеному вигляді: $M = \sum_{i=1}^n \mu_i / d_i$. Символ «/» у виразі (3.2) означає умовне прикріплення функції належності $\mu_M(d_i)$ до елемента d_i . Знак $\sum$ символічно вказує на сукупність пар $\mu_M(d_i)$ та d_i .

Остаточню ФН виступають ідентифікатором вхідних значень лінгвістичних змінних у нечіткому форматі, тобто множині значень змінної d ставляться у відповідність функції належності $\mu(d)$. Допустимі значення ЛЗ звичайно задають терм-множиною, тоді як окремий елемент цього набору іменується термом. Отже, для лінгвістичної змінної «джерело інформації» надійність визначається термами «низька», «середня», «висока», що утворюватимуть терм-множину значень.

Загалом процес формування показника, що характеризує рівень

достовірності текстових інформаційних повідомлень, можна представити як процедуру Q та набором вхідних s_i ($i = \overline{1,8}$) змінних. Якщо вказані змінні мають кількісне значення, припускається наявність діапазону їх існування, який прийнято описувати, використовуючи нижню та верхню межі значень нечітких змінних [50,76]:

$$|s_i, \overline{s_i}|, i = \overline{1, n}; |Q, \overline{Q}|. \quad (3.3)$$

Оскільки вхідна множина за суттю дослідження містить якісні змінні, тобто фактори процесу формування показника достовірності текстових інформаційних повідомлень (ДТП), необхідно встановити (можливо експертним способом) множину та межі існування значень

$$D = \{d^{(1)}, d^{(2)}, \dots, d^{(j)}\}, \quad (3.4)$$

де $d^{(k)}$, $k = \overline{1, j}$ – сукупність умовних одиниць, що можуть мати кількісну або якісну природу (за умови їх наявності), характеризується певною потужністю, яка визначається відповідним індексним показником j .

Логічно стверджувати, що вихідна змінна Q з межами визначення із (3.3) матиме подання деякою множиною з елементами в умовних одиницях

$$Q = \{q^{(1)}, q^{(2)}, \dots, q^{(g)}\}. \quad (3.5)$$

Універсальні множини (3.3) – (3.5) визначають області визначення вхідних та вихідних ЛЗ (у нашому випадку – факторів, що впливають на якість процесу формування показника ДТП) і забезпечують реалізацію залежності (3.1). При цьому лінгвістичні змінні одночасно оцінюються за допомогою якісних

лінгвістичних термів, таких як «низький», «середній», «високий», «малий», «великий» тощо, тобто термінами звичайної мови.

Формалізоване представлення бази нечітких знань, яка виконує фундаментальну функцію в структурі нечітко-логічного висновку, може бути реалізоване у вигляді матриці, що систематизує відповідні знання в узагальненій формі. [45,47]. Вона встановлює зв'язок між вхідними змінними або факторами, що впливають на процес, та його результатом – вихідною змінною. Проектування матриці знань здійснено з використанням множини нечітких логічних висловлювань, сформульованих за правилами типу «якщо-тоді-інакше», «якщо-або-тоді-інакше», «якщо-і-тоді», основу яких заклав італійський вчений Ібрагім Мамдані, що запропонував власний метод нечіткого логічного виведення, відомий як метод Мамдані [78]. На основі вказаної матриці (логічних висловлювань або логічного виведення) формуються нечіткі логічні рівняння, що обумовлюють отримання числових величин функцій належності (ФН) та інтегральний прогноз показника рівня достовірності текстових інформаційних повідомлень.

Слід зауважити, що в нечітких рівняннях операція $\vee$ вказує на отримання $\max$, а операція $\wedge$ відповідно $\min$. Це означає, що для двох значень функцій належності μ_1 та μ_2 матимемо такі комбінації отримання результату:

$$\mu_1 \vee \mu_2 = \max(\mu_1, \mu_2) = \begin{cases} \mu_1, \text{ якщо } \mu_1 \geq \mu_2, \\ \mu_2, \text{ якщо } \mu_1 < \mu_2, \end{cases} \quad (3.6)$$

$$\mu_1 \wedge \mu_2 = \min(\mu_1, \mu_2) = \begin{cases} \mu_1, \text{ якщо } \mu_1 \leq \mu_2, \\ \mu_2, \text{ якщо } \mu_1 > \mu_2. \end{cases} \quad (3.7)$$

У системах нечіткої логіки специфічні завдання, пов'язані з обробкою нечіткої інформації вирішуються в ході реалізації трьох основних етапів – фазифікації, виведення і дефазифікації. Узагальнена суть етапів полягає в тому, що у результаті фазифікації чіткі вхідні дані представлені у вигляді нечітких

величин, етап виведення забезпечує опрацювання нечітких даних за допомогою правил та отримання нечітких результатів, дефазифікація завершує процес використання нечіткої логіки при дослідженні та вирішенні конкретної проблематики перетворенням нечітких результатів у чіткі вихідні дані [79,80].

Наведені концептуальні положення у застосуванні до моделювання конкретної системи прогностичного оцінювання показника рівня достовірності текстових інформаційних повідомлень на базі нечіткої логіки зводяться до постановки та розв'язання сформульованих нижче завдань [81,82].

Фазифікація.

На початковому етапі виконується виокремлення та групування за функціональними ознаками лінгвістичних змінних – основних чинників процесу формування показника рівня достовірності текстових повідомлень.

У процесі моделювання було сформовано універсальну терм-множину значень, яка охоплює відповідні лінгвістичні терми, що відповідають визначеним лінгвістичним змінним. На цій основі реалізовано побудову ієрархічної моделі логічного виведення, структура якої узгоджується з багаторівневим представленням факторів впливу та їх описових оцінок. Найвищий рівень цієї моделі відповідає за визначення прогнозованого рівня достовірності текстових повідомлень, який подано у вигляді нечіткої множини.

Подальше чисельне опрацювання охоплює нормування функцій належності, що асоційовані з лінгвістичними змінними, на основі відповідних матриць попарного порівняння. Ці функції відображають міру приналежності лінгвістичних термів до визначених інтервальних квантів універсальної множини. Візуалізація функцій належності здійснюється шляхом побудови суміщених графічних залежностей для кожного терму, що забезпечує наочне представлення розподілу їх значень.

У структурі моделі знань реалізовано нечіткі логічні правила типу «якщо — то», які формалізують залежність між вхідними оцінками та вихідним прогнозом рівня достовірності текстових повідомлень. Це дозволяє створити

узгоджену систему логічних рівнянь, що базується на функціях належності та узагальненій матриці знань, і забезпечує аналітичне відтворення зв'язку між вхідними та вихідними параметрами моделі.

На завершальному етапі виконується **дефазифікація**, що передбачає обчислення числового значення рівня достовірності на основі методу центра ваги (або центра мас) плоскої фігури, обмеженої графіком відповідної функції належності та віссю абсцис. У цьому контексті використовуються функції належності лінгвістичних змінних, які визначені в межах універсальної множини, що окреслює допустимий простір оцінювання.

3.2. Проектування універсальної терм-множини значень

Створення універсальних терм-множин значень є актуальним питанням, оскільки дозволяє оптимізувати опрацювання нечітких даних за умов великої кількості параметрів, а також сприяє уніфікації моделей і методів та вважається важливим етапом при створенні ефективних нечітких систем. Від правильної розробки термів залежить точність, інтерпретація та зручність роботи з такими системами. Так у задачах встановлення адекватності інформаційних повідомлень терм-множина повинна бути достатньо точною, щоб врахувати всі можливі варіанти впливу на отримувані текстові повідомлення, і водночас інтуїтивно зрозумілою для користувача. Незважаючи на широкий спектр застосувань, процес проектування терм-множин залишається складним завданням, що потребує врахування багатьох факторів: особливостей вхідних даних; цільової галузі застосування; вимог до точності й адаптивності системи; розуміння потреб користувачів інтернет-новин.

Разом із тим, створення терм-множини значень ЛЗ є нелінійним і багатофакторним процесом, що залежить від таких аспектів [50,75]:

- *Структура вхідних даних.* Важливо врахувати діапазон можливих значень змінної, їх частотний розподіл та особливості, що впливають на кінцевий результат.

- *Галузевий контекст.* Термінологія, яка використовується, має бути зрозумілою для конкретної сфери застосування.
- *Кількість та взаємне розташування термів.* Надмірна кількість термів може ускладнити використання системи, тоді як їх недостатня кількість знижує точність.
- *Функції належності.* Параметри функцій, що описують терми, мають бути визначені так, щоб забезпечити оптимальний баланс між точністю та адаптивністю.

Окрім технічних аспектів, важливим є також врахування людського фактору, адже нечіткі системи часто мають суб'єктивне орієнтування. Тому проектування термів має бути орієнтованим на розуміння та зручність для кінцевого користувача.

Для досягнення поставлених цілей дослідження виконано попереднє структурування процесу формування терм-множини значень, розглядаючи його крізь призму ієрархічних складових процедур, ідентифікованих за допомогою певних груп лінгвістичних змінних. Формування складових груп ґрунтується на функціональному визначенні лінгвістичних змінних та характерові їх впливу на кінцевий результат. У цьому контексті виділено три ключові складові груп ЛЗ:, організаційна, авторська та користувацька.

У підсумку лінгвістична змінна, що ідентифікує інтегральний показник рівня достовірності інформаційних повідомлень, подається у вигляді функції

$$Q = F_Q(B, T, L), \quad (3.1)$$

аргументами якої виступають лінгвістичні змінні другого рівня B, T, L . На вказаному рівні лінгвістична змінна B орієнтована на чинники організаційного спрямування; лінгвістична змінна T визначає процедуру-функцію, аргументами якої виступають ЛЗ, що стосуються автора і контенту повідомлення; лінгвістична змінна L характеризує частку показника достовірності текстових

інформаційних повідомлень, що оцінюється відношенням (лояльністю) користувачів стосовно отриманих текстових повідомлень [45,74].

Сутність лінгвістичних змінних другого рівня відтворено формалізованими виразами, що описують їх функціональну залежність від ЛЗ третього рівня ієрархії.

$$B = F_B(b_1, b_2, b_3), \quad (3.2)$$

де: b_1 – ЛЗ «джерело інформації»; b_2 – ЛЗ «перевірка фактів»; b_3 – ЛЗ «багаторазове публікування» (перевірка через пошук багаторазового публікування).

$$T = F_T(t_1, t_2, t_3), \quad (3.3)$$

де: t_1 – ЛЗ «професіоналізм автора»; t_2 – ЛЗ «об'єктивність автора»; t_3 – ЛЗ «інформативність контенту».

$$L = F_L(l_1, l_2), \quad (3.4)$$

де: l_1 – ЛЗ «спростування і критика»; l_2 – ЛЗ «соціальна довіра».

Проведене структурування вихідної бази даних зумовило перехід до наступного етапу, який передбачає формування універсальної бази або терм-множини значень функцій лінгвістичних змінних. Нагадаємо, що терм-множина слугує інструментом для відображення реальних (найчастіше числових) або умовних, за їхньої відсутності, меж дії чи впливу, які здійснюються конкретною лінгвістичною змінною на процес. Лінгвістичні терми надають описову інформацію про вагомість лінгвістичних змінних, співвіднесену до вузлів інтервального поділу універсальної множини.

Для впорядкування відповідних даних та забезпечення їх структурованого подання розроблено таблицю 3.1, у якій узагальнено ключові характеристики лінгвістичних змінних: символічне позначення, вербальну інтерпретацію,

відповідну терм-множину як універсальну область значень, а також визначену множину лінгвістичних термів. Останні згруповано у форматі тривимірних шкал, що забезпечує здійснення якісної інтерпретації можливих станів лінгвістичних змінних в межах установлених інтервальних підмножин універсальної множини.

Таблиця 3.1

Універсальні множини термів лінгвістичних змінних

Змінна	Лінгвістичний опис змінної	Універсальна база значень (множина U)	Лінгвістичні терми (терм-множина H)
b_1	Джерело інформації (надійність)	(1-5) у. о.	Низька, середня, висока
b_2	Перевірка фактів	(1-5) у. о.	Нечаста, часта, постійна
b_3	Багаторазове публікування (кількість)	(1-5) у. о.	Мала, середня, велика
t_1	Професіоналізм автора	(1-5) у. о.	Низький, середній, високий
t_2	Об'єктивність автора	(1-5) у. о.	Низька, середня, висока
t_3	Інформативність контенту повідомлення	(1-5) у. о.	Низька, середня, висока
l_1	Спростування і критика (частота)	(1-5) у. о.	Мала, середня, значна
l_2	Соціальна довіра	(1-5) у. о.	Низька, середня, висока

Універсальні множини термів лінгвістичних змінних у процесі оцінювання рівня достовірності інформаційних повідомлень підтвердили ефективність застосування структурованого підходу. Запропонована терм-множина уможливорює якісний опис меж впливу лінгвістичних змінних, відображаючи їхні значення як у числовій, так і в умовній формах. Це, у свою чергу, забезпечує універсальність моделі та підвищує точність оцінювання достовірності інформаційних повідомлень.

Групування термів у тривимірні шкали сприяє детальному аналізу станів змінних у квантах універсальної множини, що підсилює можливості моделі для адаптації до різних контекстів застосування. Застосування такого підходу

створює передумови для подальшого вдосконалення методів оцінювання достовірності інформації в умовах невизначеності

Результати формування універсальної множини термів лінгвістичних змінних для визначення рівня достовірності інформаційних повідомлень підтверджують ефективність запропонованої моделі. Використання термів у тривимірних шкалах дозволяє точно оцінювати вплив нечітких змінних, що підвищує точність аналізу достовірності текстової інформації та забезпечує адаптивність до різних умов.

3.3. Метод нечіткого логічного виведення як основа формування показника рівня достовірності повідомлень

У сучасних умовах інформаційного середовища, що характеризується високою швидкістю поширення даних та їх великою кількістю традиційні методи аналізу, засновані на чітких логічних принципах, не завжди здатні адекватно відобразити складність та неоднозначність реальних інформаційних процесів. У цьому контексті метод нечіткого логічного виведення пропонує ефективний підхід, здатний враховувати невизначеність і варіативність даних. Метод базується на принципах нечіткої логіки, яка дозволяє працювати з неповною або нечіткою інформацією, що є характерною для реальних ситуацій. Він стає теоретичною основою для формування показника рівня достовірності повідомлень і дає змогу більш точно оцінити ймовірність істинності інформації в умовах невизначеності. Даний підхід дозволяє інтегрувати різні джерела інформації та обробляти їх з урахуванням нечітких параметрів, що значно покращує точність і достовірність результатів [74].

На основі виконаного структурування за функціональними складовими та рівнями вагомості факторів процесу формування показника рівня достовірності інформаційних повідомлень розроблено метод логічного виведення, графічна модель якого відтворює багаторівневу ієрархію зв'язків між компонентами

вихідної бази даних та визначає алгоритм розрахунку значень функцій належності лінгвістичних змінних з врахуванням лінгвістичних термів. Модель у вигляді багаторівневої графічної структури втілює алгоритмічний механізм, який поетапно формує ієрархію визначення показника рівня достовірності текстових інформаційних повідомлень, починаючи з лінгвістичних змінних базового рівня (рисунок 3.1). Для ідентифікації цих змінних відповідно до принципів нечіткої логіки було сформовано терм-множину значень ЛЗ, що забезпечує адекватне відображення невизначеності та нечіткості факторів. Застосування такого підходу охоплює широкий спектр галузей — від систем штучного інтелекту до управління якістю, що свідчить про його універсальність

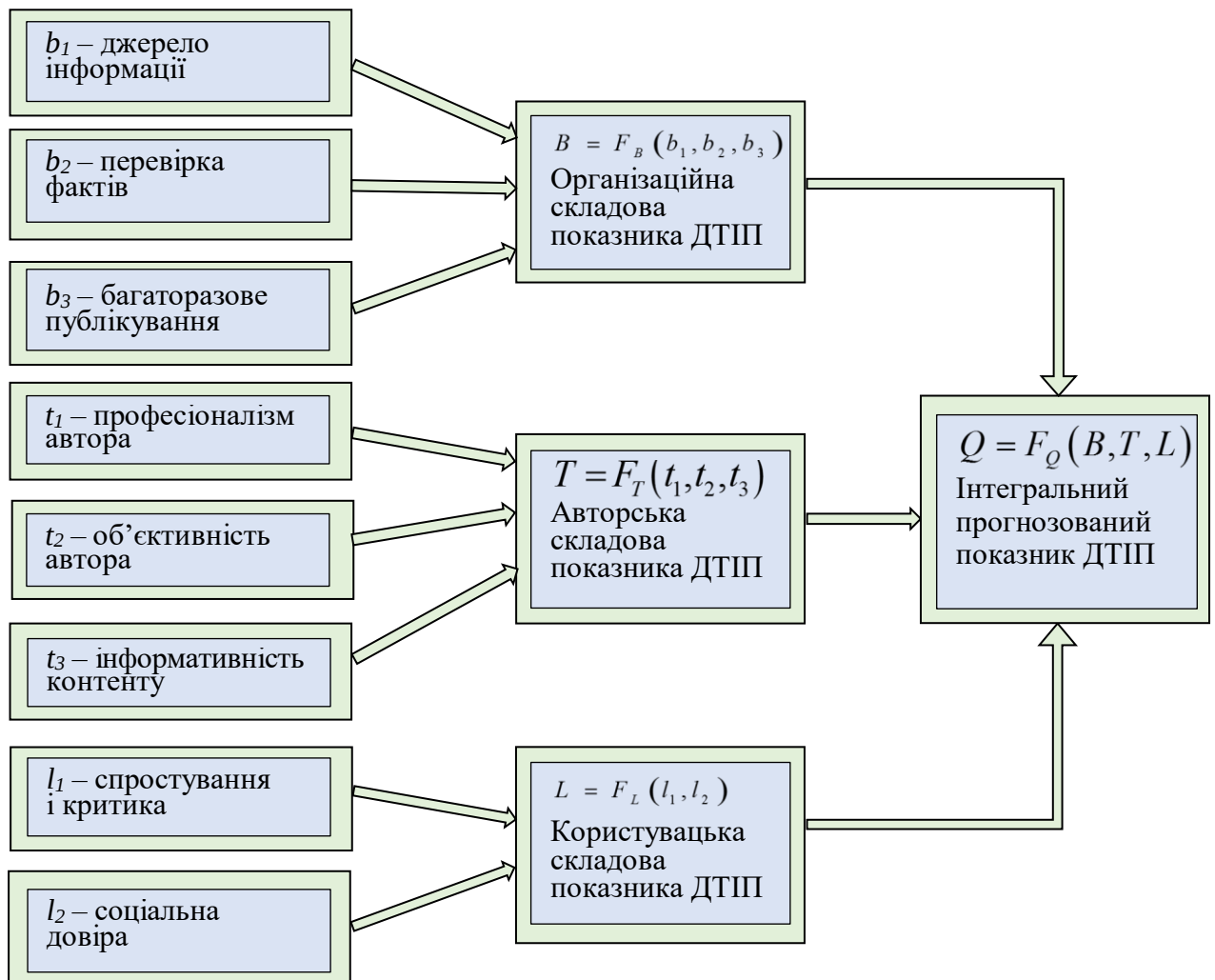


Рис. 3.1. Модель методу логічного виведення: формування показника достовірності текстових інформаційних повідомлень

Багаторівнева модель побудована таким чином, щоб забезпечити логіку

формування індикатора достовірності текстових повідомлень за принципом «від окремого до загального». Цей підхід враховує часткові показники достовірності, які визначаються на кожному рівні ієрархії. Вплив лінгвістичних змінних найнижчого рівня оцінюється за допомогою універсального набору значень термів і відповідних лінгвістичних характеристик змінних, що дозволяє забезпечити системний аналіз текстової інформації.

Такий підхід до оцінювання достовірності інформації зумовив врахування не тільки локальних, але й загальних факторів впливу на правдивість текстових інформаційних повідомлень, розміщених на нижчих рівнях ієрархії. На вищих рівнях отримані результати поетапно інтегруються у загальний прогностичний показник достовірності. Універсальна терм-множина забезпечує стандартизований опис значень лінгвістичних змінних, а також підвищує точність і гнучкість моделі в процесі опрацювання текстових даних. Таким чином, багаторівнева модель сприяє не лише точному та неупередженому прогностичному оцінюванню достовірності текстових повідомлень, але й адаптації у певній мірі до різних контекстів та умов використання.

3.4. Функції належності лінгвістичних змінних як чинників формування показника оцінювання рівня достовірності текстових повідомлень

З урахуванням наведених вище передумов, у першу чергу сформованої терм-множини значень лінгвістичних змінних, лінгвістичний терм «показник рівня достовірності текстових інформаційних повідомлень» Q відображено у вигляді нечіткої множини [74]

$$Q = \left\{ \frac{\mu_q(u_1)}{u_1}, \frac{\mu_q(u_2)}{u_2}, \dots, \frac{\mu_q(u_n)}{u_n} \right\}, \quad (3.5)$$

де: $Q \subset U$; $\mu_q(u_i)$ – функція належності елемента $u_i \in U$ до множини Q .

У формулі (3.5) функції належності подано у вигляді нормалізованих

значень, які застосовуються під час побудови та розв'язання нечітких логічних рівнянь. Зв'язок між функціями належності лінгвістичних змінних і ранжуванням відповідних лінгвістичних термів цих змінних визначено через наступні співвідношення:

$$\frac{\mu_1}{r_1} = \frac{\mu_2}{r_2} = \dots = \frac{\mu_n}{r_n}, \quad (3.6)$$

де: $\mu_i = \mu_q(u_i)$; $r_i = r_q(u_i)$ для всіх $i=1, \dots, n$; r – ранги лінгвістичних змінних.

Нормовані значення функцій належності повинні задовольняти умову нормування, тобто: $\mu_1 + \mu_2 + \dots + \mu_n = 1$.

Ранги лінгвістичних змінних сформовано на основі експертного оцінювання, яке відображається через функції належності і реалізується поетапно, частину яких виконано вище. Короткий виклад вказаних процедур наведено нижче.

Визначення лінгвістичних змінних та їх термів. Лінгвістичні змінні та віднесені до них словесні терми, що описують їх значення, наведені у табл. 3.1.

Визначення шкали оцінювання. Вибирається шкала для порівняння (наприклад, 1–9 за методом Сааті, що відповідає задачі). Кожному терму (рангу) лінгвістичної змінної призначається відповідне числове значення, що дозволяє виконувати математичні операції.

Формування функцій належності. Для кожного терму створюється функція належності, яка визначає, як об'єкт належить певному терму. У дисертаційному дослідженні, форма функцій належності гаусівська.

Експертне оцінювання. Експерти визначають відносну важливість критеріїв або альтернатив у парі. Ці оцінки задаються через лінгвістичні змінні, які пізніше трансформуються у числові значення.

Перетворення лінгвістичних оцінок у числові значення. На основі функцій належності кожна лінгвістична оцінка замінюється нечітким числом, що враховує невизначеність і суб'єктивність.

Формування матриці попарних порівнянь. Отримані числові значення

достовірності текстових інформаційних повідомлень, за умови позитивних значень лінгвістичних термів універсальної бази U та максимумів функцій належності у точках квантового розподілу множини U [74]. Математичне моделювання означеної задачі наведено у наступній інтерпретації:

$$\left. \begin{aligned} Q = F(b_k, t_j, l_p) &\rightarrow \max, k = \overline{1,3}; j = \overline{1,3}; \\ b_k > 0, t_j > 0, l_p > 0, p &= \overline{1,2}; \\ \mu_q(u_i) &\rightarrow \max, u_i \in U, Q \subset U, i = \overline{1,5}. \end{aligned} \right\}. \quad (3.9)$$

Формалізований запис завдання (3.9) містить компоненти універсальної терм-множини значень b_k, t_j, l_p стосовно лінгвістичних змінних процесу.

Остаточно для кожної лінгвістичної змінної та відповідних їй трьох лінгвістичних термів множини H побудовано квадратні обернено симетричні матриці $A = a_{ij}$ ($a_{ij} = r_i / r_j$; $i, j = \overline{1,5}$), елементи яких згенеровано за допомогою порівняння рангів між собою. Загальний вигляд матриці наведено нижче [50].

$$A = \begin{bmatrix} 1 & \frac{r_2}{r_1} & \frac{r_3}{r_1} & \frac{r_4}{r_1} & \frac{r_5}{r_1} \\ \frac{r_1}{r_2} & 1 & \frac{r_3}{r_2} & \frac{r_4}{r_2} & \frac{r_5}{r_2} \\ \dots & \dots & \dots & \dots & \dots \\ \frac{r_1}{r_5} & \frac{r_2}{r_5} & \frac{r_3}{r_5} & \frac{r_4}{r_5} & 1 \end{bmatrix}. \quad (3.10)$$

Застосування теоретичних положень, викладених вище, здійснимо почергово для всіх чинників процесу формування показника достовірності текстових інформаційних повідомлень [74].

Лінгвістична змінна b_1 «джерело інформації» (надійність).

Універсальна база значень ЛЗ знаходиться в інтервалі (1-5) умовних одиниць в опорних точках $U(b_1) = [1; 2; 3; 4; 5]$ і характеризується рівнями

«низька», «середня» та «висока» відповідно до таблиці 1. Значення змінної обираються в межах інтервалу $(1 \div 9)$, з урахуванням того, що для забезпечення психологічної межі при порівнянні об'єктів використовується дев'ять градацій відмінностей між ними [46]. Відзначимо, що у випадку узгодженості експертних оцінок під час порівнянь вказується лише останній рядок матриці, тоді як інші її елементи визначаються на основі залежності $a_{ji} = a_{5j} / a_{5i}; i, j = \overline{1,5}$. Найнижчий ранг ЛЗ у точці u_1 для терму «низька» встановлено цифрою 1. На інтервалі ранг змінної зростає.

Матриця попарних порівнянь для терму «низька» має такий вигляд:

$$A_{\text{низька}}(b_1) = \begin{bmatrix} 1 & 8/9 & 5/9 & 3/9 & 1/9 \\ 9/8 & 1 & 5/8 & 3/8 & 1/8 \\ 9/5 & 8/5 & 1 & 3/5 & 1/5 \\ 9/3 & 8/3 & 5/3 & 1 & 1/3 \\ 9 & 8 & 5 & 3 & 1 \end{bmatrix}. \quad (3.11)$$

На підставі виразів (3.8) та з урахуванням елементів матриці виконано розрахунок значень функцій належності в усіх точках множини.

$$\mu_{\text{низька}}(u_1) = \frac{1}{1 + \frac{8}{9} + \frac{5}{9} + \frac{3}{9} + \frac{1}{9}} = \frac{9}{26} = 0.35; \quad \mu_{\text{низька}}(u_2) = \frac{8}{26} = 0.31.$$

$$\mu_{\text{низька}}(u_3) = \frac{5}{26} = 0.19; \quad \mu_{\text{низька}}(u_4) = \frac{3}{26} = 0.12; \quad \mu_{\text{низька}}(u_5) = \frac{1}{26} = 0.04.$$

Для терму «середня» матриця отримала таке відображення:

$$A_{\text{середня}}(b_1) = \begin{bmatrix} 1 & 6 & 9 & 5 & 1 \\ 1/6 & 1 & 9/6 & 5 & 1/6 \\ 1/9 & 6/9 & 1 & 5/9 & 1/9 \\ 1/5 & 6/5 & 9/5 & 1 & 1/5 \\ 1 & 6 & 9 & 5 & 1 \end{bmatrix}. \quad (3.12)$$

З урахуванням матриці розраховано такі значення функцій належності:

$$\mu_{\text{середня}}(u_1) = 0.05; \mu_{\text{середня}}(u_2) = 0.27; \mu_{\text{середня}}(u_3) = 0.41; \\ \mu_{\text{середня}}(u_4) = 0.23; \mu_{\text{середня}}(u_5) = 0.05.$$

Аналогічно для лінгвістичного терму «висока» сформовано матрицю

$$A_{\text{висока}}(b_1) = \begin{bmatrix} 1 & 3 & 5 & 7 & 9 \\ 1/3 & 1 & 5/3 & 7/3 & 9/3 \\ 1/5 & 3/5 & 1 & 7/5 & 9/5 \\ 1/7 & 3/7 & 5/7 & 1 & 9/7 \\ 1/9 & 3/9 & 5/9 & 7/9 & 1 \end{bmatrix}. \quad (3.13)$$

Отримано такі числові значення для функцій належності:

$$\mu_{\text{висока}}(u_1) = 0.04; \mu_{\text{висока}}(u_2) = 0.12; \mu_{\text{висока}}(u_3) = 0.20; \\ \mu_{\text{висока}}(u_4) = 0.28; \mu_{\text{висока}}(u_5) = 0.36.$$

Наступний крок – нормування значень функцій належності, отриманих на основі матриць попарних порівнянь рангів лінгвістичних змінних, необхідне для забезпечення коректності та узгодженості результатів аналізу. Основні причини нормування зумовлені такими його особливостями [50,76].

Узгодженість значень функцій належності. Значення функцій належності повинні лежати в межах $[0, 1]$, що відповідає базовим принципам нечіткої логіки. Нормування дозволяє привести обчислені значення до цього діапазону, зберігаючи пропорційність між ними.

Порівнюваність між критеріями. Нормування робить можливим коректне порівняння значень функцій належності різних лінгвістичних змінних, які

можуть мати різні масштаби. Це особливо важливо, коли дані використовуються в складних моделях, таких як прийняття рішень або оцінювання альтернатив.

Запобігання спотворенням результатів. Без нормування значення можуть бути некоректно інтерпретовані або неправильно враховані в подальших розрахунках, таких як агрегування нечіткої інформації або дефазифікація.

Забезпечення сумісності з математичними моделями. У багатьох алгоритмах нечіткої логіки, таких як нечіткий висновок або кластеризація, використовуються нормалізовані функції належності. Ненормовані значення можуть призводити до математичних і логічних помилок у моделі.

Виконано нормування функцій належності множенням їх значень на коефіцієнт $k_r = 1/\max \mu_r(u_i)$, ($i=1, \dots, 5$), де r означає «низька», «середня», «висока». На практиці це ділення кожного значення функцій належності на максимальне значення з відповідного набору

За результатами використання формули $\mu_{r_n}(d_i) = k_r \times \mu_r(u_i)$ для лінгвістичної змінної «джерело інформації» (надійність джерела) при визначених термах «низька», «середня», «висока» отримано наступні нормовані значення функцій належності:

$$\begin{aligned} \mu_{\text{низька}_n}(u_1) &= 1; \mu_{\text{низька}_n}(u_2) = 0.86; \mu_{\text{низька}_n}(u_3) = 0.54; \\ \mu_{\text{низька}_n}(u_4) &= 0.34; \mu_{\text{низька}_n}(u_5) = 0.11; \\ \mu_{\text{середня}_n}(u_1) &= 0.12; \mu_{\text{середня}_n}(u_2) = 0.66; \mu_{\text{середня}_n}(u_3) = 1; \\ \mu_{\text{середня}_n}(u_4) &= 0.56; \mu_{\text{середня}_n}(u_5) = 0.12; \\ \mu_{\text{висока}_n}(u_1) &= 0.11; \mu_{\text{середня}_n}(u_2) = 0.33; \mu_{\text{висока}_n}(u_3) = 0.66; \\ \mu_{\text{висока}_n}(u_4) &= 0.78; \mu_{\text{висока}_n}(u_5) = 1. \end{aligned}$$

Числові пріоритети функцій належності лінгвістичної змінної «джерело інформації» (надійність) для термів «низька», «середня», «висока» представлено у вигляді нормованих значень відповідно до узагальненого виразу (3.5) у рамках нечітких множин:

$$\begin{aligned}
\text{надійність низька} &= \left\{ \frac{1}{1}; \frac{0,86}{2}; \frac{0,54}{3}; \frac{0,34}{4}; \frac{0,11}{5} \right\} \text{ у. о.}; \\
\text{надійність середня} &= \left\{ \frac{0,12}{1}; \frac{0,66}{2}; \frac{1}{3}; \frac{0,56}{4}; \frac{0,12}{5} \right\} \text{ у. о.}; \\
\text{надійність висока} &= \left\{ \frac{0,11}{1}; \frac{0,33}{2}; \frac{0,66}{3}; \frac{0,78}{4}; \frac{1}{5} \right\} \text{ у. о.}
\end{aligned} \tag{3.14}$$

Важливим етапом у використанні інструментів нечіткої логіки для аналізу процесів із наявністю невизначеності є візуалізація функцій належності [80]. Основна мета графічного представлення функцій належності для лінгвістичної змінної «джерело інформації» (надійність) за термами «низька», «середня», «висока» полягає у демонстрації рівня відповідності джерела інформації конкретному рівню надійності залежно від вибраного терму. Така візуалізація підкреслює, що надійність не є категорично визначеною характеристикою (наприклад, «надійне» чи «ненадійне»), а змінюється поступово. Значення функцій належності ілюструють ступінь відповідності ЛЗ кожному терму та допомагають оцінити рівень надійності джерела інформації.

Однією з ключових властивостей візуалізації в контексті аналізу нечітких множин є здатність до часткового перекривання функцій належності. Така особливість дозволяє здійснювати адаптивне моделювання зв'язків між елементами досліджуваної системи. У розглянутому варіанті це означає, що окреме джерело інформації може одночасно належати кільком термам із різним ступенем відповідності. Завдяки цьому забезпечено розширений підхід до оцінки надійності, що сприяє уникненню жорстких меж між термами та дозволяє більш адаптивно відображати рівень достовірності інформаційних повідомлень. У рамках дослідження використано Гаусові функції належності, які мають низку суттєвих переваг порівняно з трикутними та трапецієподібними функціями. Таке рішення є особливо актуальним для моделювання процесів, у яких характер змін не є різким, а відбувається поступово, що часто спостерігається у сфері оцінювання довіри до інформації.

На рис. 3.2 наведено графічне представлення нечітких множин, що

характеризують лінгвістичну змінну «джерело інформації». У даній візуалізації відображено функції належності для термів «низька», «середня» та «висока», а також подано відповідні нормовані функції належності, виражені рівнянням (3.14). Такий підхід дав змогу отримати цілісне уявлення про градацію надійності джерел інформації та їхній розподіл у просторі значень. Деталізація функцій належності дозволила точніше описати рівень довіри до конкретного джерела та встановити межі між різними рівнями оцінки. Це важливо для побудови моделей оцінювання достовірності інформаційних повідомлень, де враховується не лише номінальне значення параметра, а й ступінь його невизначеності.

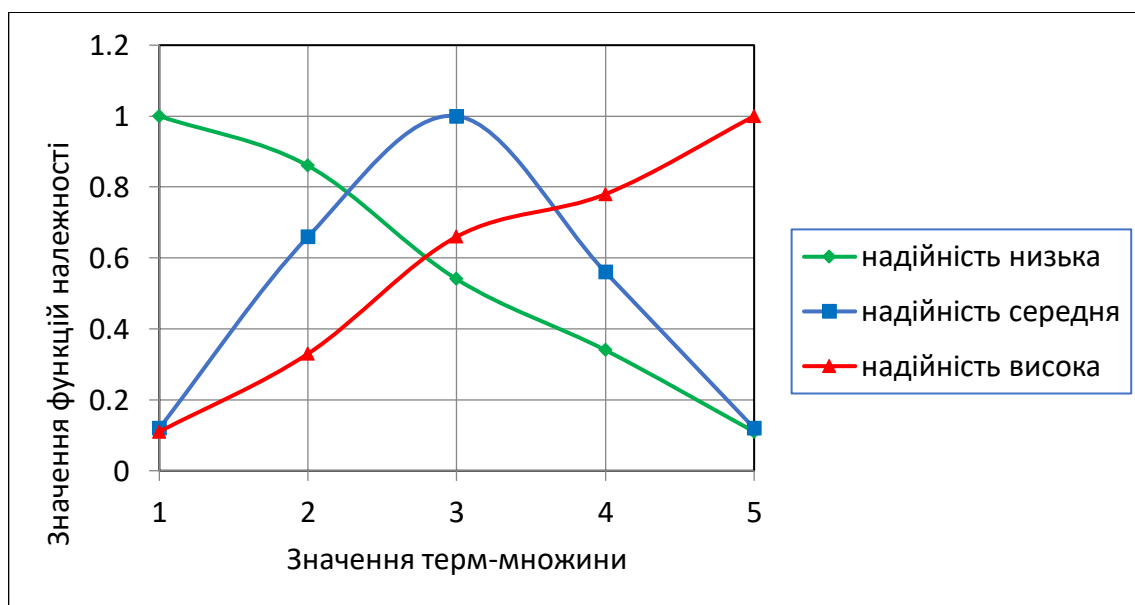


Рис. 3.2. Графічне представлення функцій належності лінгвістичної змінної «джерело інформації»

Наступна лінгвістична змінна досліджується за схемою попередньої.

Лінгвістична змінна b_2 «перевірка фактів».

Універсальна база значень ЛЗ знаходиться в інтервалі (1-5) умовних одиниць в опорних точках $U(b_2) = [1; 2; 3; 4; 5]$ і характеризується рівнями «нечаста», «часта», «постійна». Не повторюючи обґрунтувань, що стосувалися попередньої змінної, перейдемо до побудови матриць попарних порівнянь рангів

змінної, фіксування нормованих значень функцій належності та побудови відповідного графічного відображення.

Для терму «нечаста» матриця має такий вигляд.

$$A_{\text{нечаста}}(b_2) = \begin{bmatrix} 1 & 6/9 & 4/9 & 2/9 & 1/9 \\ 9/6 & 1 & 4/6 & 2/6 & 1/6 \\ 9/4 & 6/4 & 1 & 2/4 & 1/4 \\ 9/2 & 6/2 & 4/2 & 1 & 1/2 \\ 9 & 6 & 4 & 2 & 1 \end{bmatrix}. \quad (3.15)$$

Розрахунок значень функцій належності дав такі результати.

$$\mu_{\text{нечаста}}(u_1) = 0.41; \mu_{\text{нечаста}}(u_2) = 0.27; \mu_{\text{нечаста}}(u_3) = 0.18;$$

$$\mu_{\text{нечаста}}(u_4) = 0.09; \mu_{\text{нечаста}}(u_5) = 0.05.$$

Терм «часта» для ЛЗ b_2 обумовив таку матрицю.

$$A_{\text{часта}}(b_2) = \begin{bmatrix} 1 & 7 & 9 & 4 & 1 \\ 1/7 & 1 & 9/7 & 4/7 & 1/7 \\ 1/9 & 7/9 & 1 & 4/9 & 1/9 \\ 1/4 & 7/4 & 9/4 & 1 & 1/4 \\ 1 & 7 & 9 & 4 & 1 \end{bmatrix}. \quad (3.16)$$

Подібно до розрахунків, виконаних вище, отримано відповідні числові значення функцій належності ЛЗ «перевірка фактів».

$$\mu_{\text{часта}}(u_1) = 0.05; \mu_{\text{часта}}(u_2) = 0.32; \mu_{\text{часта}}(u_3) = 0.41;$$

$$\mu_{\text{часта}}(u_4) = 0.18; \mu_{\text{часта}}(u_5) = 0.05.$$

Лінгвістична змінна b_2 для терму «постійно» має подібну матрицю.

$$A_{\text{постійна}}(b_2) = \begin{bmatrix} 1 & 3 & 5 & 7 & 9 \\ 1/3 & 1 & 5/3 & 7/3 & 9/3 \\ 1/5 & 3/5 & 1 & 7/5 & 9/5 \\ 1/7 & 3/7 & 5/7 & 1 & 9/7 \\ 1/9 & 3/9 & 5/9 & 7/9 & 1 \end{bmatrix}. \quad (3.17)$$

Розрахунок за елементами матриці (3.17) та формулами (3.8) забезпечив числові значення функцій належності.

$$\mu_{\text{постійна}}(u_1) = 0.04; \mu_{\text{постійна}}(u_2) = 0.12; \mu_{\text{постійна}}(u_3) = 0.20;$$

$$\mu_{\text{постійна}}(u_4) = 0.28; \mu_{\text{постійна}}(u_5) = 0.36.$$

Нормування функцій належності для трьох термів «нечаста», «часта», «постійна» лінгвістичної змінної «перевірка фактів» дало такі результати:

$$\mu_{\text{нечаста}_n}(u_1) = 1; \mu_{\text{нечаста}_n}(u_2) = 0.66; \mu_{\text{нечаста}_n}(u_3) = 0.44;$$

$$\mu_{\text{нечаста}_n}(u_4) = 0.22; \mu_{\text{нечаста}_n}(u_5) = 0.11.$$

$$\mu_{\text{часта}_n}(u_1) = 0.11; \mu_{\text{часта}_n}(u_2) = 0.78; \mu_{\text{часта}_n}(u_3) = 1;$$

$$\mu_{\text{часта}_n}(u_4) = 0.44; \mu_{\text{часта}_n}(u_5) = 0.11.$$

$$\mu_{\text{постійна}_n}(u_1) = 0.11; \mu_{\text{постійна}_n}(u_2) = 0.33; \mu_{\text{постійна}_n}(u_3) = 0.68;$$

$$\mu_{\text{постійна}_n}(u_4) = 0.78; \mu_{\text{постійна}_n}(u_5) = 1.$$

Нормовані значення функцій належності ЛЗ «перевірка фактів» для термів «нечаста», «часта», «постійна» представлено нечіткою множиною (3.18):

$$\begin{aligned} \text{перевірка нечаста} &= \left\{ \frac{1}{1}; \frac{0,66}{2}; \frac{0,44}{3}; \frac{0,22}{4}; \frac{0,11}{5} \right\} \text{ у. о.}; \\ \text{перевірка часта} &= \left\{ \frac{0,11}{1}; \frac{0,78}{2}; \frac{1}{3}; \frac{0,44}{4}; \frac{0,11}{5} \right\} \text{ у. о.}; \\ \text{перевірка постійна} &= \left\{ \frac{0,11}{1}; \frac{0,33}{2}; \frac{0,68}{3}; \frac{0,78}{4}; \frac{1}{5} \right\} \text{ у. о.} \end{aligned} \quad (3.18)$$

Графічна візуалізація функцій належності ЛЗ наведена нижче.

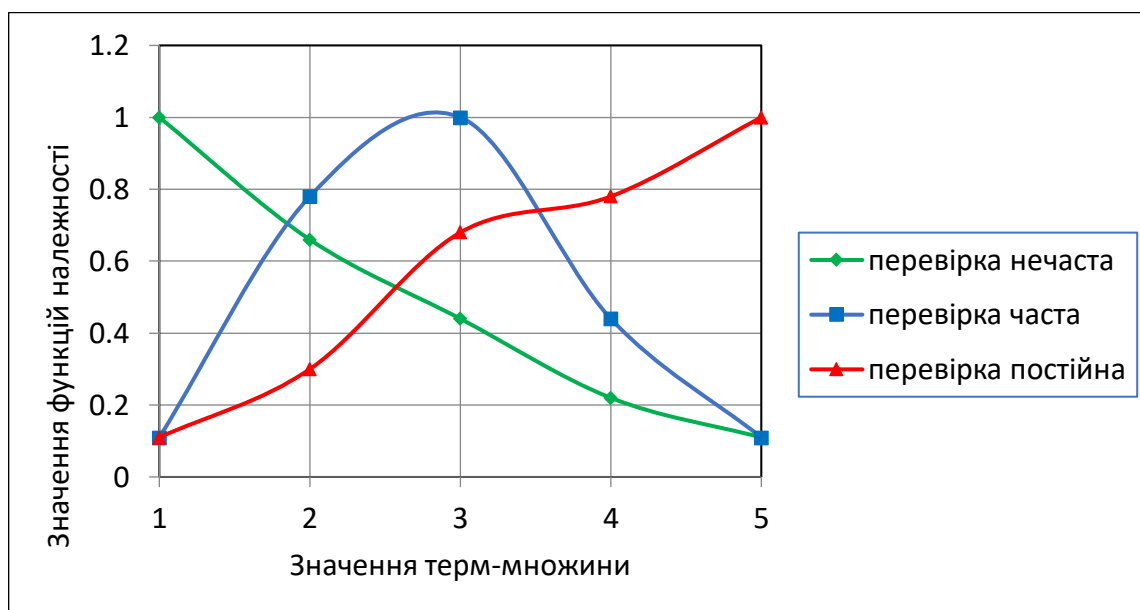


Рис. 3.3. Графічне представлення функцій належності лінгвістичної змінної «перевірка фактів»

Розрахунок та візуалізацію решти функцій належності здійснено за алгоритмом, реалізованим вище, що обумовлює мінімізацію процесу. У подальшому будемо наводити тільки результуючі нечіткі множини нормованих значень функцій належності та відповідні їм графічні відображення.

Лінгвістична змінна b_3 «багаторазове публікування» (кількість).

ЛЗ b_3 «багаторазове публікування» визначено на множині $U(b_3) = [1; 2; 3; 4; 5]$ лінгвістичними термами «мала», «середня», «велика».

Результуюча нечітка множина функцій належності:

$$\begin{aligned}
 \text{кількість мала} &= \left\{ \frac{1}{1}; \frac{0,50}{2}; \frac{0,20}{3}; \frac{0,11}{4}; \frac{0,05}{5} \right\} \text{ у. о.}; \\
 \text{кількість середня} &= \left\{ \frac{0,30}{1}; \frac{0,60}{2}; \frac{1}{3}; \frac{0,40}{4}; \frac{0,15}{5} \right\} \text{ у. о.}; \\
 \text{кількість велика} &= \left\{ \frac{0,11}{1}; \frac{0,20}{2}; \frac{0,40}{3}; \frac{0,80}{4}; \frac{1}{5} \right\} \text{ у. о.}
 \end{aligned} \tag{3.19}$$

Графічне відображення функцій належності лінгвістичної множини (3.19) наведено на рис. 3.4.

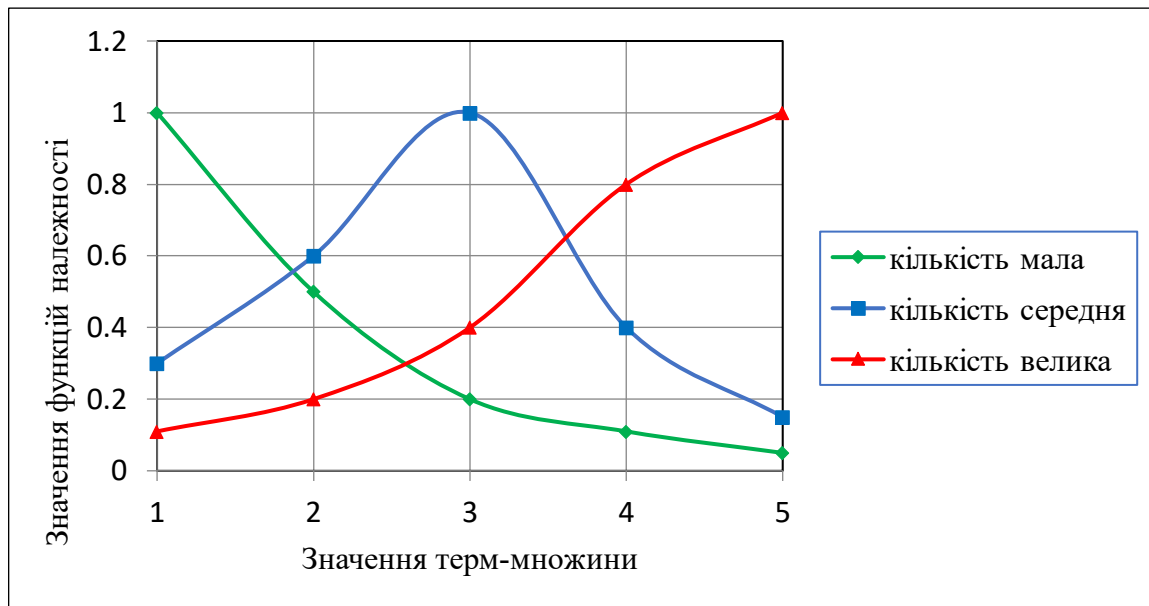


Рис. 3.4. Графічне представлення функцій належності лінгвістичної змінної «багаторазове публікування»

Переходимо до авторської складової ЛЗ: t_1 «професіоналізм автора»; t_2 «об'єктивність автора»; t_3 «інформативність контенту».

Лінгвістична змінна t_1 «професіоналізм автора».

ЛЗ t_1 «професіоналізм автора» задано в інтервалі (1-5) умовних одиниць в опорних точках множини $U(t_1) = [1; 2; 3; 4; 5]$ лінгвістичними термами «низький», «середній», «високий». Нечітка множина функцій належності лінгвістичної змінної має такий вигляд:

$$\begin{aligned}
 \text{низький} &= \left\{ \frac{1}{1}; \frac{0,80}{2}; \frac{0,40}{3}; \frac{0,30}{4}; \frac{0,10}{5} \right\} \text{ у. о.}; \\
 \text{середній} &= \left\{ \frac{0,20}{1}; \frac{0,40}{2}; \frac{1}{3}; \frac{0,50}{4}; \frac{0,15}{5} \right\} \text{ у. о.}; \\
 \text{високий} &= \left\{ \frac{0,11}{1}; \frac{0,25}{2}; \frac{0,30}{3}; \frac{0,80}{4}; \frac{1}{5} \right\} \text{ у. о.}
 \end{aligned} \tag{3.20}$$

Просторову графічну візуалізацію функцій належності лінгвістичної змінної «перевірка фактів» за нечіткими термами «низький», «середній», «високий», виражених нечіткою множиною (3.20), наведено на рис. 3.5.

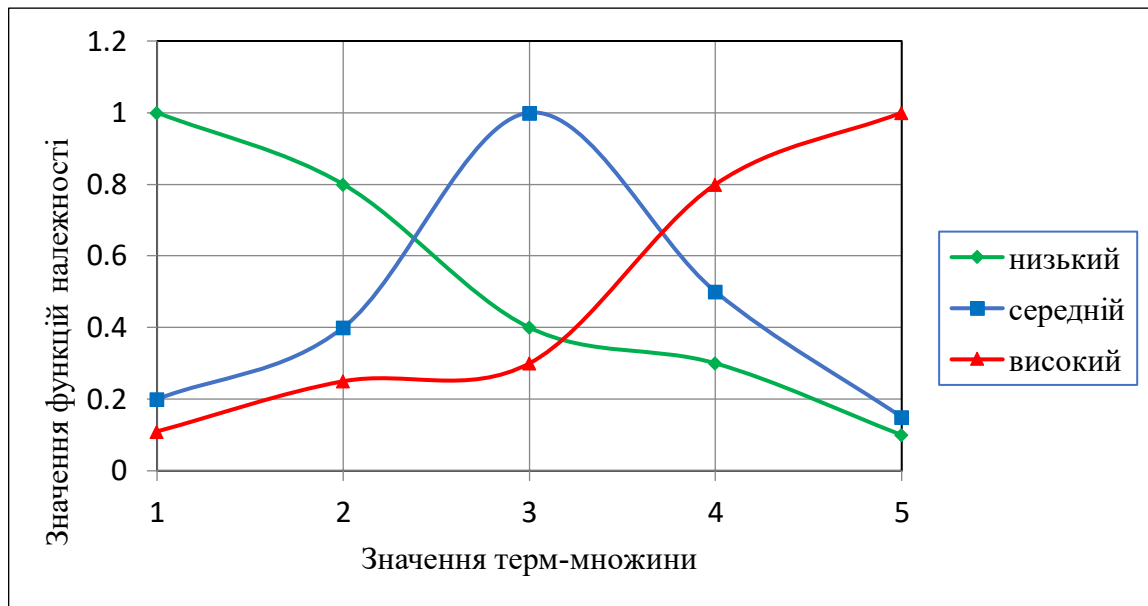


Рис. 3.5. Графічне представлення функцій належності лінгвістичної змінної «професіоналізм автора»

Лінгвістична змінна t_2 «об'єктивність автора».

Числові значення ЛЗ t_2 «об'єктивність автора» розраховано в базових точках множини $U(t_2) = [1; 2; 3; 4; 5]$ лінгвістичними термами «низька», «середня», «висока». Нечітку множину нормованих значень функцій належності ЛЗ стосовно вказаних нечітких термів відтворено виразом (3.21).

$$\begin{aligned}
 \text{низька} &= \left\{ \frac{1}{1}; \frac{0,85}{2}; \frac{0,45}{3}; \frac{0,35}{4}; \frac{0,10}{5} \right\} \text{ у. о.}; \\
 \text{середня} &= \left\{ \frac{0,15}{1}; \frac{0,44}{2}; \frac{1}{3}; \frac{0,85}{4}; \frac{0,15}{5} \right\} \text{ у. о.}; \\
 \text{висока} &= \left\{ \frac{0,06}{1}; \frac{0,12}{2}; \frac{0,5}{3}; \frac{0,60}{4}; \frac{1}{5} \right\} \text{ у. о.}
 \end{aligned} \tag{3.21}$$

Візуальне відображення нечіткої множини (3.21) для означених нечітких термів «низька», «середня», «висока» подано на рис. (3.6).

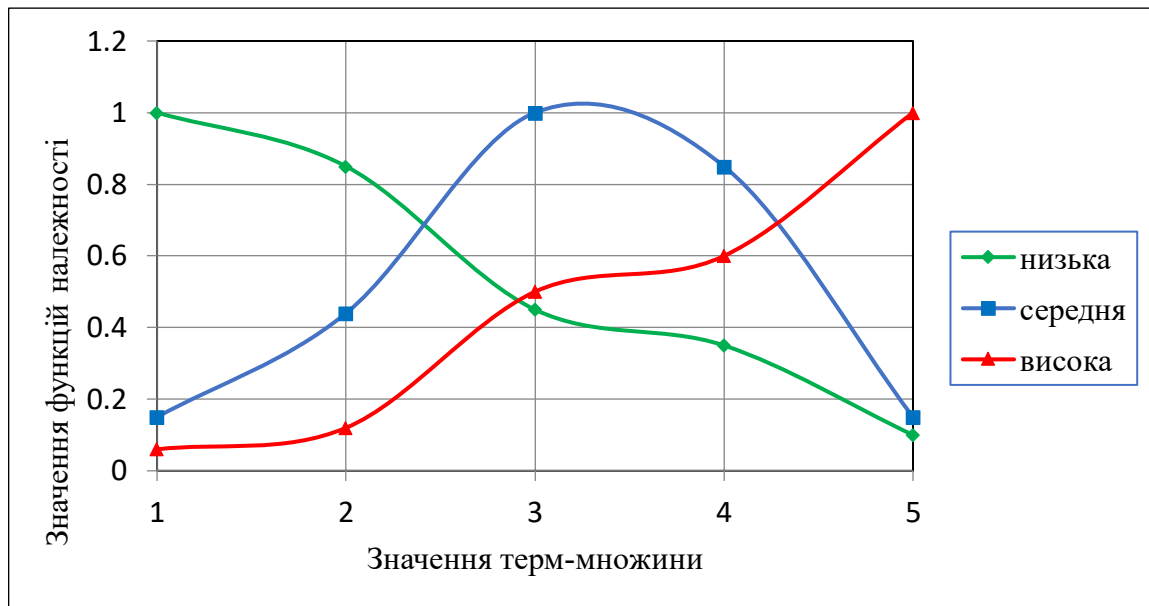


Рис. 3.6. Графічне представлення функцій належності лінгвістичної змінної «об'єктивність автора»

Лінгвістична змінна t_3 «інформативність контенту повідомлення».

ЛЗ t_3 «інформативність контенту повідомлення» ідентифіковано значеннями функцій належності в точках базової множини $U(t_3) = [1; 2; 3; 4; 5]$. Нечітка терм множина утворена термами «низька», «середня», «висока». Нечітку множину нормованих значень ФН оформлено виразом (3.22).

$$\begin{aligned}
 \text{низька} &= \left\{ \frac{1}{1}; \frac{0,55}{2}; \frac{0,33}{3}; \frac{0,22}{4}; \frac{0,11}{5} \right\} \text{ у. о.}; \\
 \text{середня} &= \left\{ \frac{0,11}{1}; \frac{0,89}{2}; \frac{1}{3}; \frac{0,33}{4}; \frac{0,11}{5} \right\} \text{ у. о.}; \\
 \text{висока} &= \left\{ \frac{0,11}{1}; \frac{0,15}{2}; \frac{0,44}{3}; \frac{0,78}{4}; \frac{1}{5} \right\} \text{ у. о.}
 \end{aligned} \tag{3.22}$$

Графічну візуалізацію функцій належності лінгвістичної змінної «інформативність контенту повідомлення» наведено на рисунку 3.7. Вказані ФН відповідають нечітким термам «низька», «середня» та «висока» і визначені нечіткою множиною (3.22). Функції належності відображають ступінь належності кожного значення змінної до відповідного терму, що дозволяє моделювати невизначеність та суб'єктивність процесу оцінювання інформативності контенту – важливої складової достовірності повідомлень.

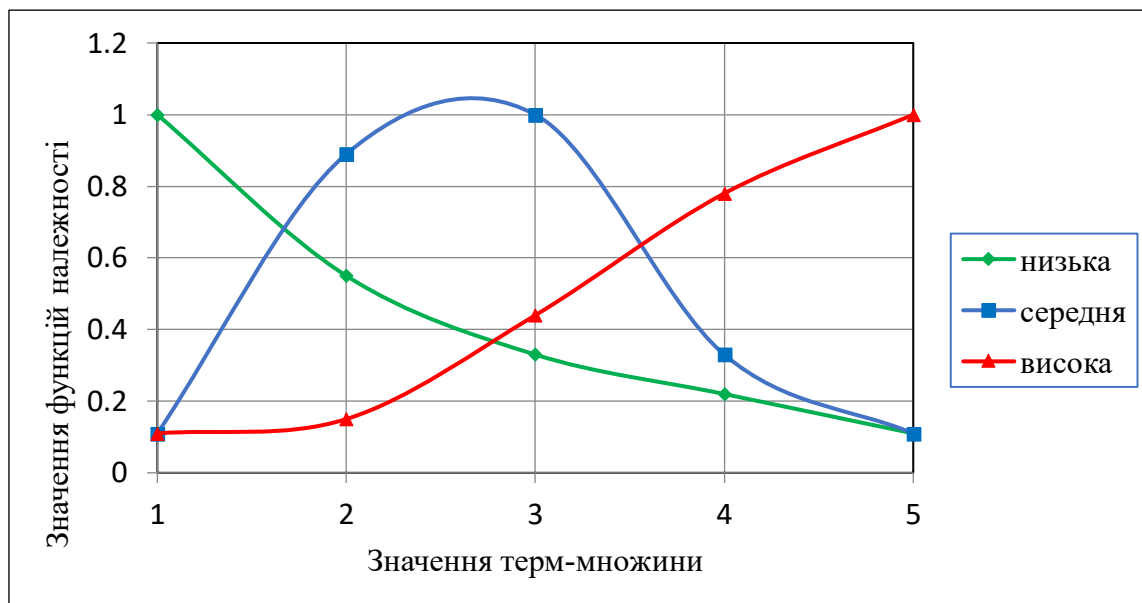


Рис. 3.7. Графічне представлення функцій належності лінгвістичної змінної «інформативність контенту повідомлення»

Завершує етап дослідження функцій належності користувацька складова, частка якої у ході формування ДТІП визначена лінгвістичними змінними l_1 «спростування і критика» (частота) і l_2 «соціальна довіра», числовими значеннями ФН, узгодженими з відповідними нечіткими термами.

Лінгвістична змінна l_1 «спростування і критика» (частота).

Лінгвістичну змінну l_1 «спростування і критика» (частота) задано нормованими числовими значеннями ФН в точках множини $U(l_1) = [1; 2; 3; 4; 5]$ для нечітких термів «мала», «середня», «значна». Нечітку множину для зазначених термів представлено виразом (3.23), який формалізує функції належності для моделювання невизначеності та нечіткості в різних системах.

$$\begin{aligned}
 \text{частота мала} &= \left\{ \frac{1}{1}; \frac{0,30}{2}; \frac{0,22}{3}; \frac{0,10}{4}; \frac{0,05}{5} \right\} \text{ у. о.}; \\
 \text{частота середня} &= \left\{ \frac{0,25}{1}; \frac{0,60}{2}; \frac{1}{3}; \frac{0,30}{4}; \frac{0,25}{5} \right\} \text{ у. о.}; \\
 \text{частота значна} &= \left\{ \frac{0,11}{1}; \frac{0,40}{2}; \frac{0,50}{3}; \frac{0,87}{4}; \frac{1}{5} \right\} \text{ у. о.}
 \end{aligned} \tag{3.23}$$

Візуалізацію нечіткої множини (3.23) наведено на рис. 3.8.

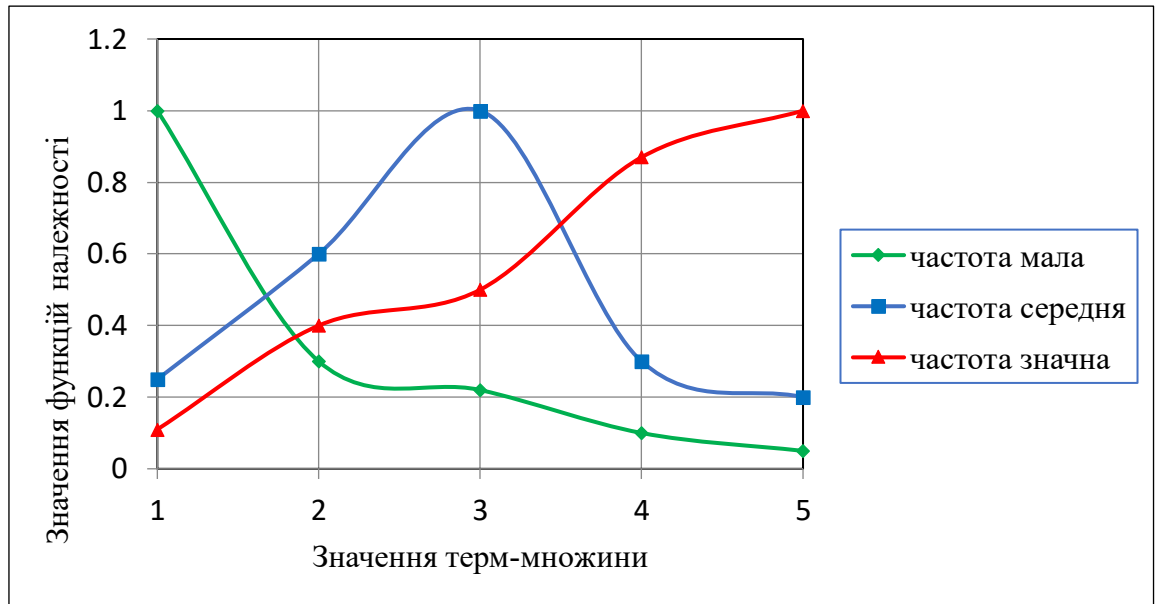


Рис. 3.8 Графічне представлення функцій належності лінгвістичної змінної «спростування і критика» (частота)

Оскільки вказана лінгвістична змінна стосується ролі користувачів, наведемо декілька думок стосовно рис. 3.8. Візуалізація показує, як часто інформаційні повідомлення піддаються критиці або спростуванню у відповідних термах. Результат може допомогти зрозуміти, наскільки поширеними є інформаційні повідомлення, що викликають критику. Графік може виявити, як часто з'являються повідомлення, що викликають критику.

Лінгвістична змінна l_2 «соціальна довіра».

Лінгвістичну змінну l_2 «соціальна довіра» описано значеннями функцій належності в точках множини $U(l_1) = [1; 2; 3; 4; 5]$ для нечітких термів «низька», «середня», «висока». Нечітку множину задано виразом (3.24).

$$\begin{aligned}
 \text{низька} &= \left\{ \frac{1}{1}; \frac{0,75}{2}; \frac{0,45}{3}; \frac{0,35}{4}; \frac{0,10}{5} \right\} \text{ у. о.}; \\
 \text{середня} &= \left\{ \frac{0,15}{1}; \frac{0,44}{2}; \frac{1}{3}; \frac{0,85}{4}; \frac{0,15}{5} \right\} \text{ у. о.}; \\
 \text{висока} &= \left\{ \frac{0,06}{1}; \frac{0,12}{2}; \frac{0,5}{3}; \frac{0,60}{4}; \frac{1}{5} \right\} \text{ у. о.}
 \end{aligned} \tag{3.24}$$

Графічну візуалізацію лінгвістичної змінної наведено на рис. 3.9.

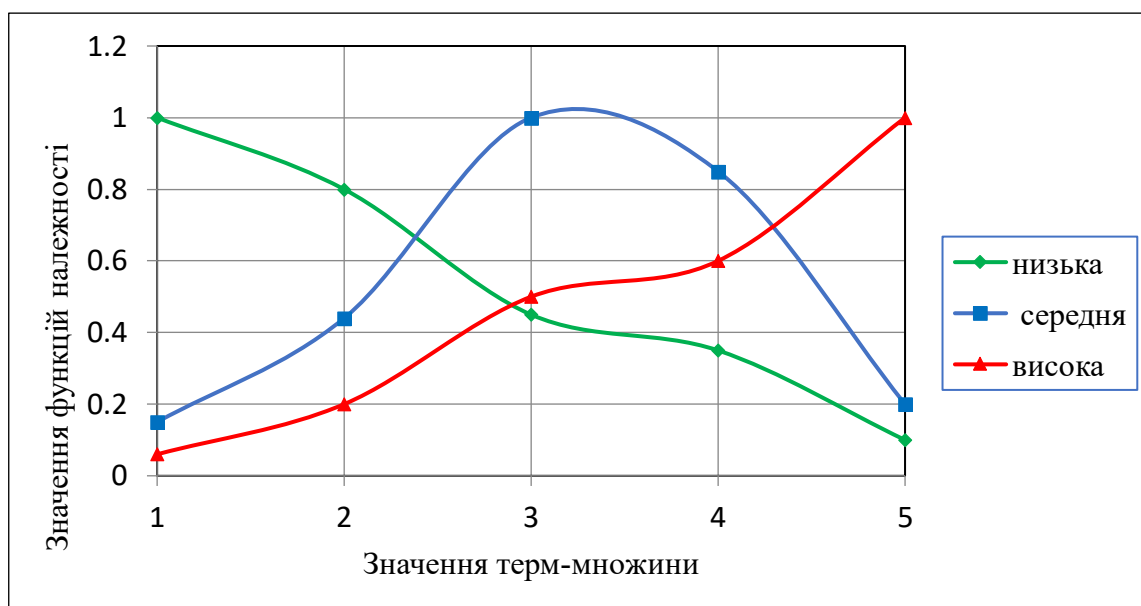


Рис. 3.9. Графічне представлення функцій належності лінгвістичної змінної «соціальна довіра»

Результати візуалізації здатні забезпечити обґрунтоване прийняття рішень, спрямованих на підвищення соціальної довіри. Наприклад, низький рівень може вимагати заходів для зміцнення комунікації між владою і громадянами. Низька соціальна довіра населення може бути наслідком корупції, економічної кризи чи поганих комунікацій уряду. Домінування високої соціальної довіри може свідчити про ефективну політику, сильне громадянське суспільство або високу якість соціальних послуг. Графічна візуалізація ЛЗ «соціальна довіра» дозволяє оцінити поточний стан, тенденції та визначити можливі напрями для покращення соціальних взаємодій.

Графічне відображення функцій належності сприяє аналізу динаміки взаємодії елементів нечітких множин на заданому рівні дослідження та дозволяє ідентифікувати базу даних, пов'язану з формуванням показника достовірності текстової інформації. Окрім того, отримані оцінки функціональної належності лінгвістичних змінних слугують основою для створення нечіткої бази знань і системи нечітких логічних рівнянь, що використовуються при проектуванні та обчисленні інтегрального показника достовірності інформаційних повідомлень.

3.5. Проектування нечітких матриць знань та нечітких логічних рівнянь

В основі досліджуваної проблематики лежить прогнозне оцінювання достовірності текстових повідомлень, для якого традиційні методи часто виявляються недостатніми через неможливість повного врахування невизначеності та неоднозначності чинників впливу. Вирішальне значення при цьому має взаємозв'язок факторного аналізу ключових чинників достовірності з методами нечіткої логіки, що забезпечує можливість формування адаптивних моделей. Такий підхід дозволяє враховувати широкий спектр характеристик, які не піддаються однозначній формалізації, але істотно впливають на майбутню оцінку достовірності. Використання засобів нечіткої логіки в поєднанні з факторним аналізом створює умови для підвищення ефективності прогнозування рівня достовірності інформаційних повідомлень на основі компонент нечіткої бази знань – матриць знань та нечітких логічних рівнянь.

Згідно з теоретичними положеннями нечіткої логіки [50,76] *нечітка база знань* – це широке поняття, яке містить усю необхідну інформацію для роботи нечіткої системи. Основне призначення бази знань – забезпечення нечіткої системи інформацією, потрібною для аналізу даних і прийняття рішень. Вона стосується: опису параметрів процесу через *лінгвістичні змінні* та нечіткі термножини; *функцій належності*, що виражаються через математичні вирази, які визначають, наскільки елемент належить до певної нечіткої множини; *нечіткого наслідування*, як сукупності умовно-логічних залежностей у форматі «якщо-тоді»; *матриці знань* – це частина бази знань, яка має більш вузьку функцію і формується таблицею або структурою, що визначає залежності між входами та виходами системи. У матриці зберігаються узагальнені правила у вигляді значень або індексів, які відповідають результатам застосування нечітких логічних правил.

Отже, база знань є ширшою концепцією, яка охоплює всі аспекти знань у нечіткій системі, тоді як матриця знань є інструментом для організації правил і

зручного відображення їхніх взаємозв'язків.

З огляду на відзначене та враховуючи напрацювання попередніх параграфів стосовно перших двох складових нечіткої бази знань, виконаємо проектування нечітких матриць знань через формалізоване відтворення нечіткої дедуктивної процедури та формування нечітких логічних рівнянь, які дозволяють ефективно оцінювати достовірність текстових повідомлень. Вони мають на меті забезпечити надійну основу для подальшого розвитку інтелектуальних систем аналізу тексту в різних галузях, включаючи журналістику, освіту, маркетинг та державне управління.

Процес формування показника достовірності текстових інформаційних повідомлень у межах моделі нечіткої дедуктивної процедури рис. 3.1 узагальнює результати оцінювання якості у групах лінгвістичних змінних, що відповідає принципам та практичним підходам до аналізу і обробки текстових даних. Оскільки досліджувані процеси реалізуються із застосуванням методів нечіткої логіки, завершальні етапи, пов'язані з визначенням інтегрального показника достовірності даних, передбачають побудову нечітких логічних висловлювань для лінгвістичних змінних усіх рівнів моделі рис. 3.1, що стали основою проектування матриць знань та відповідних їм систем нечітких логічних рівнянь. Структура матриць знань повинна враховувати експертні оцінки, які визначають сукупний вплив лінгвістичних змінних на процес оцінювання достовірності текстових інформаційних повідомлень залежно від комбінацій їхніх лінгвістичних термів [74, 81, 82].

Отже, для функції найвищого рівня, позначеного лінгвістичною змінною Q і відношенням (3.1), яка визначає інтегральний показник прогнозованої достовірності повідомлень, задаємо множину нечітких лінгвістичних термів $H(Q) = \langle \text{низька, середня, висока} \rangle$.

Наступний рівень складають: лінгвістична змінна B , орієнтована на чинники організаційного спрямування з термами $H(B) = \langle \text{низька, середня, висока} \rangle$; лінгвістична змінна T , що визначає авторську частку і контекст

повідомлення, з термами $H(T) = \langle \text{низька, середня, висока} \rangle$; лінгвістична змінна L , що характеризує відношення (лояльність) користувачів, та оцінюється термами $H(L) = \langle \text{низька, середня, висока} \rangle$.

З урахуванням наведених даних нечітке логічне виведення для найвищого рівня має такий формалізований вигляд:

ЯКЩО $(B = \text{низька}) \text{ І } (B = \text{середня}) \text{ І } (B = \text{висока})$

$\text{І } (T = \text{низька}) \text{ І } (T = \text{середня}) \text{ І } (T = \text{висока})$

$\text{І } (L = \text{низька}) \text{ І } (L = \text{середня}) \text{ І } (L = \text{висока}),$

ТОДІ $(Q = \text{низька}) \text{ І } (Q = \text{середня}) \text{ І } (Q = \text{висока}).$

Сформульоване логічне виведення обумовлює розроблення відповідної матриці знань у вигляді таблиці 3.2. На основі матриці знань таблиці 3.2 виконано побудову нечітких логічних рівнянь, що дозволило отримати значення функцій належності для термів, представлених у четвертому стовпці таблиці. Ці рівняння описують логічні залежності між входами системи та їх відповідністю лінгвістичним термам.

Таблиця 3.2

Матриця знань для лінгвістичної змінної Q

Організаційна складова ДТП B	Авторська складова ДТП T	Користувацька складова ДТП L	Інтегральний показник ДТП Q
низька	низька	низька	низька
низька	середня	низька	
середня	низька	середня	середня
середня	середня	висока	
висока	висока	середня	висока
висока	висока	висока	

Терм «низька» спричиняє отримання такого нечіткого рівняння:

$$\mu_{\text{низька}}(Q) = \mu_{\text{низька}}(B) \wedge \mu_{\text{низька}}(T) \wedge \mu_{\text{низька}}(L) \vee \mu_{\text{низька}}(B) \wedge \mu_{\text{низька}}(T) \wedge \mu_{\text{низька}}(L).$$

Відповідне рівняння забезпечує функцію належності з термом «середня»:

$$\mu_{\text{середня}}(Q) = \mu_{\text{середня}}(B) \wedge \mu_{\text{низька}}(T) \wedge \mu_{\text{середня}}(L) \vee \mu_{\text{середня}}(B) \wedge \mu_{\text{середня}}(T) \wedge \mu_{\text{висока}}(L).$$

Опис функції належності з термом «висока» обумовлює таке рівняння:

$$\mu_{\text{висока}}(Q) = \mu_{\text{висока}}(B) \wedge \mu_{\text{висока}}(T) \wedge \mu_{\text{середня}}(L) \vee \mu_{\text{висока}}(B) \wedge \mu_{\text{висока}}(T) \wedge \mu_{\text{висока}}(L).$$

Застосуємо вказаний підхід до лінгвістичних змінних наступного рівня ієрархії. При цьому потрібно враховувати, що ЛЗ цього рівня потребують врахування нечітких термів змінних, що ідентифіковані у таблиці 3.1 термножиною H . Таким чином, лінгвістична змінна B , орієнтована на чинники організаційного спрямування, обумовила такий варіант логічного виведення:

ЯКЩО $(b_1) = (\text{низька, середня, висока})$

$I(b_2) = (\text{нечаста, часта, постійна})$

$I(b_3) = (\text{мала, середня, велика})$

ТОДІ $(B) = (\text{низька, середня, висока})$.

Таблиця 3.3

Матриця знань для лінгвістичної змінної B

Джерело інформації b_1 (надійність)	Перевірка фактів b_2	Багаторазове публікування (кількість) b_3	Організацій на складова ДТІП В
низька	нечаста	мала	низька
середня	нечаста	мала	
середня	часта	велика	середня
висока	часта	мала	
висока	постійна	середня	висока
середня	постійна	велика	

Нечіткі логічні рівняння для обчислення значень функцій належності лінгвістичної змінної B за означеними в матриці термами подано нижче.

Для терму «низька» логічне рівняння має такий вигляд:

$$\mu_{\text{низька}}(B) = \mu_{\text{низька}}(b_1) \wedge \mu_{\text{нечаста}}(b_2) \wedge \mu_{\text{мала}}(b_3) \vee \mu_{\text{середня}}(b_1) \wedge \mu_{\text{нечаста}}(b_2) \wedge \mu_{\text{мала}}(b_3).$$

Терм «середня» обумовить наступне рівняння:

$$\mu_{\text{середня}}(B) = \mu_{\text{середня}}(b_1) \wedge \mu_{\text{часта}}(b_2) \wedge \mu_{\text{велика}}(b_3) \vee \\ \vee \mu_{\text{висока}}(b_1) \wedge \mu_{\text{часта}}(b_2) \wedge \mu_{\text{мала}}(b_3).$$

Терму «висока» відповідатиме функція належності:

$$\mu_{\text{висока}}(B) = \mu_{\text{висока}}(b_1) \wedge \mu_{\text{постійна}}(b_2) \wedge \mu_{\text{середня}}(b_3) \vee \\ \vee \mu_{\text{середня}}(b_1) \wedge \mu_{\text{постійна}}(b_2) \wedge \mu_{\text{велика}}(b_3).$$

Логічне виведення для лінгвістична змінної T , що визначає авторську частку і контекст повідомлення, для термів низька, середня, висока таке:

ЯКЩО $(t_1) = (\text{низький, середній, високий})$

$I(t_2) = (\text{низька, середня, висока})$

$I(t_3) = (\text{низька, середня, висока}),$

ТОДІ $(T) = (\text{низька, середня, висока}).$

Матрицю знань згідно логічного виведення подана нижче.

Таблиця 3.4

Матриця знань для лінгвістичної змінної T

Професіоналізм автора t_1	Об'єктивність автора t_2	Інформативність контенту t_3	Технічна складова ДТІП T
низький	висока	низька	низька
низький	низька	середня	
середній	середня	середня	середня
середній	висока	низька	
високий	середня	висока	висока
високий	висока	висока	

Функції належності ЛЗ T «Технічна складова» для множини термів «низька, середня, висока» отримано за системою логічних рівнянь, наведених нижче.

Для терму «низька» нечітке рівняння матиме вигляд:

$$\mu_{\text{низька}}(T) = \mu_{\text{низький}}(t_1) \wedge \mu_{\text{висока}}(t_2) \wedge \mu_{\text{низька}}(t_3) \vee \\ \vee \mu_{\text{низький}}(t_1) \wedge \mu_{\text{низька}}(t_2) \wedge \mu_{\text{середня}}(t_3).$$

Терму «середня» відповідає таке нечітке логічне рівняння:

$$\mu_{\text{середня}}(T) = \mu_{\text{середній}}(t_1) \wedge \mu_{\text{середня}}(t_2) \wedge \mu_{\text{середня}}(t_3) \vee \\ \vee \mu_{\text{середній}}(t_1) \wedge \mu_{\text{висока}}(t_2) \wedge \mu_{\text{низька}}(t_3).$$

Логічне рівняння, що відноситься до терму «висока»:

$$\mu_{\text{висока}}(T) = \mu_{\text{високий}}(t_1) \wedge \mu_{\text{середня}}(t_2) \wedge \mu_{\text{висока}}(t_3) \vee \\ \vee \mu_{\text{високий}}(t_1) \wedge \mu_{\text{висока}}(t_2) \wedge \mu_{\text{висока}}(t_3).$$

Логічне висловлювання для лінгвістичної змінної L , що характеризує користувацьку складову з термами низька, середня, висока:

ЯКЩО $(l_1) = (\text{мала, середня, значна})$

І $(l_2) = (\text{низька, середня, висока})$

ТОДІ $(L) = (\text{низька, середня, висока})$.

Для наведеного логічного висловлювання побудовано матрицю знань.

Таблиця 3.5

Матриця знань для лінгвістичної змінної L

Спростування і критика (частота) l_1	Соціальна довіра l_2	Користувацька складова ДТП L
мала	середня	низька
середня	низька	
значна	низька	середня
середня	середня	
значна	середня	висока
значна	висока	

Систему нечітких рівнянь для підрахунку значення функції належності змінної L наведено для множини термів «низька, середня, висока».

Рівняння для терму «низька» матиме вигляд:

$$\mu_{\text{низька}}(L) = \mu_{\text{мала}}(l_1) \wedge \mu_{\text{середня}}(l_2) \vee \mu_{\text{середня}}(l_1) \wedge \mu_{\text{низька}}(l_2).$$

Терм «середня» обумовлює таке рівняння:

$$\mu_{\text{середня}}(L) = \mu_{\text{значна}}(l_1) \wedge \mu_{\text{низька}}(l_2) \vee \mu_{\text{середня}}(l_1) \wedge \mu_{\text{середня}}(l_2).$$

Логічне рівняння, що відповідає терму «висока»:

$$\mu_{\text{висока}}(L) = \mu_{\text{значна}}(l_1) \wedge \mu_{\text{середня}}(l_2) \vee \mu_{\text{значна}}(l_1) \wedge \mu_{\text{висока}}(l_2).$$

Виконані вище дослідження лінгвістичних змінних, дотичних до процесу оцінювання достовірності інформаційних повідомлень, обумовили отримання функцій належності і побудову нечітких логічних рівнянь, що уможливило розрахунок інтегрального показника оцінювання рівня ДТІП. Застосування методів нечіткої логіки сприяє підвищенню точності оцінювання та дозволяє коригувати параметри в реальному часі, що підвищує гнучкість методики і забезпечує вдосконалення алгоритмів обробки текстових даних.

Висновки до розділу 3

1. Охарактеризовано ключові поняття теорії нечітких множин, зокрема: лінгвістичні змінні, функції належності, терм-множини значень, лінгвістичні терми факторів, бази нечітких знань, матриці знань, нечіткі логічні рівняння, а також експертні логічні висновки.

2. Виконано структурування процесу формування показника рівня достовірності текстових інформаційних повідомлень через створення трьох основних складових: організаційної, технічної та користувацької, які представлені відповідними групами лінгвістичних змінних.

3 Сформовано інформаційну базу даних, яка встановлює зв'язок між змінними, їхньою лінгвістичною природою, діапазонами значень універсальної терм-множини та визначеними лінгвістичними термами, заданими нечіткою тривимірною шкалою для відображення якісних характеристик змінних.

4. Розроблено багаторівневу модель логічного виведення, яка відображає логічні процеси побудови показника рівня достовірності текстових інформаційних повідомлень з урахуванням груп лінгвістичних змінних.

5. Розроблено та досліджено матриці попарних порівнянь для лінгвістичних термів лінгвістичних змінних у п'яти точках поділу універсальної

множини. На основі компонентів власних векторів матриць визначено ранги термів для обчислення значень функцій належності.

6. Виконано графічне візуальне відображення функцій належності для кожного терму лінгвістичних змінних із використанням функцій Гауса. Їх застосування дозволяє досягти плавного переходу між значеннями, що є перевагою порівняно з трикутними та трапецієвидними функціями. Такий підхід є особливо доцільним для моделювання процесів із поступовими змінами.

7. Запроектовано нечіткі матриці знань, у яких відображено формалізовану реалізацію нечіткого логічного висновку, а також сформовано нечіткі логічні рівняння, виконання яких обумовить розрахунок показника достовірності текстових інформаційних повідомлень.

РОЗДІЛ 4. ІНФОРМАЦІЙНА ТЕХНОЛОГІЯ ОЦІНЮВАННЯ ДОСТОВІРНОСТІ ТЕКСТОВИХ ПОВІДОМЛЕНЬ

У сучасних умовах глобалізації інформаційного простору, стрімкого зростання обсягів текстових повідомлень і загроз масового поширення дезінформації особливої актуальності набуває розроблення ефективних технологій для оцінювання достовірності інформації ще до її споживання кінцевим користувачем. Надійна і своєчасна верифікація контенту майбутніх повідомлень дозволяє запобігти негативному впливу фейкових повідомлень на громадську думку, знизити ризики маніпуляцій та підвищити інформаційну безпеку як на індивідуальному, так і на суспільному рівні.

Поточний розділ присвячено розробленню інформаційної технології прогностного оцінювання достовірності текстових інформаційних повідомлень, яка базується на поєднанні аналітичних, обчислювальних та логіко-семантичних підходів. На початковому етапі виконано експериментальний розрахунок інтегрального показника прогностного рівня достовірності текстових повідомлень, що передбачає формування прикладної бази даних, в основі якої таблиці значень функцій належності лінгвістичних змінних відповідно до якісних градацій терм-множин, розраховані в попередньому розділі. Визначення кількісного показника прогностної достовірності текстових інформаційних повідомлень виконано за допомогою методу центру мас.

У подальшому акцент зроблено на програмному моделюванні процесу оцінювання ДТІП із використанням методів багатокритеріальної оптимізації. Такий підхід дає змогу враховувати множину факторів впливу, формалізувати процедуру прийняття рішень і забезпечити адаптацію системи до різних сценаріїв застосування.

Наступним етапом розділу є функціональне моделювання процесу створення інформаційної технології, яке реалізовано з використанням методології IDEF0. Це дозволило відобразити логіку побудови технології у вигляді структурованих взаємозалежних процесів, уточнити рольових учасників,

інформаційні потоки, вхідні й вихідні параметри, а також ресурси, необхідні для реалізації кожного з етапів.

Сукупність наведених прикладних аспектів сформувало цілісну концепцію інформаційної технології, спрямованої на забезпечення високого рівня надійності оцінювання достовірності текстових повідомлень і підвищення стійкості інформаційного середовища до маніпулятивних впливів.

4.1. Експериментальний розрахунок показника достовірності текстових інформаційних повідомлень

Результати обробки лінгвістичних змінних B, T, L другого рівня ієрархії моделі логічного виведення обумовили створення нечіткої множини, яка описує лінгвістичну змінну Q «показник рівня достовірності текстових інформаційних повідомлень» з функціями належності, визначеними для трьох нечітких термів: «низька», «середня», «висока». Кількісні значення цієї змінної визначено у трьох точках поділу універсальної множини $\{q_1, q_2, q_3\}$. У підсумку лінгвістичний терм Q «показник достовірності текстових інформаційних повідомлень» представлено у вигляді нечіткої множини [46,74], елементи якої сформовано у вигляді сукупності пар:

$$Q(B, T, L) = \left\{ \frac{\mu_{\text{низька}}(Q)}{q_1}, \frac{\mu_{\text{середня}}(Q)}{q_2}, \frac{\mu_{\text{висока}}(Q)}{q_3} \right\}, \quad (4.1)$$

де q_1, q_2, q_3 – визначають кількісні значення лінгвістичної змінної Q стосовно вказаних вище термів у точках поділу універсальної множини.

Відповідно до запропонованої моделі логічного виведення прогнозування результатів для будь-яких значень вхідних параметрів здійснено на основі нечіткого логічного висновку, побудованого на експертно сформованій базі знань, представлений правилами типу «ЯКЩО – ТОДІ», матрицею знань для певної лінгвістичної змінної (див. табл. 3.2) і нечіткою множиною (4.1).

На початковому етапі моделювання для використаних лінгвістичних змінних заплановано сформувати відповідні таблиці з нормалізованими значеннями функцій належності у п'яти контрольних точках розподілу універсальної множини значень, що охоплюють характерні області кожного терм-множини. Такий підхід дозволяє забезпечити достатню точність у подальшому числовому представленні нечітких понять і дає змогу відобразити поступовий перехід між значеннями, властивий лінгвістичним змінним. Зазначені таблиці виступають базовим інструментом для інтерпретації семантики лінгвістичних оцінок, оскільки дають змогу визначити числові значення функцій належності для кожного з відповідних термів.

Нормалізовані дані становлять основу для кількісної обробки вхідної інформації, оскільки їх використання дозволяє узгодити різномірні показники та забезпечити коректність подальших обчислювальних процедур. Зокрема, вони були застосовані у розрахунках на основі визначених нечітких логічних рівнянь, що дало змогу здійснити не лише формальне опрацювання числових значень, а й надати математичну інтерпретацію лінгвістичних характеристик.

Таблиця 4.1

Функції належності терм-множини $H(b_1)$
(джерело інформації – надійність)

u_i , у. о.	1	2	3	4	5
$\mu_{\text{низька}}(u_i)$	1	0,86	0,54	0,34	0,11
$\mu_{\text{середня}}(u_i)$	0,12	0,66	1	0,56	0,12
$\mu_{\text{висока}}(u_i)$	0,11	0,33	0,66	0,78	1

Таблиця 4.2

Функції належності терм-множини $H(b_2)$
(перевірка фактів)

u_i , у. о.	1	2	3	4	5
$\mu_{\text{нечаста}}(u_i)$	1	0,66	0,44	0,22	0,11
$\mu_{\text{часта}}(u_i)$	0,11	0,78	1	0,44	0,11
$\mu_{\text{постійна}}(u_i)$	0,11	0,33	0,68	0,78	1

Таблиця 4.3

Функції належності терм-множини $H(b_3)$
(багаторазове публікування – кількість)

u_i , у. о.	1	2	3	4	5
$\mu_{мала}(u_i)$	1	0,50	0,20	0,11	0,05
$\mu_{середня}(u_i)$	0,30	0,60	1	0,40	0,15
$\mu_{велика}(u_i)$	0,11	0,20	0,40	0,80	1

Таблиця 4.4

Функції належності терм-множини $H(t_1)$
(професіоналізм автора)

u_i , у. о.	1	2	3	4	5
$\mu_{низький}(u_i)$	1	0,80	0,40	0,30	0,10
$\mu_{середній}(u_i)$	0,20	0,40	1	0,50	0,15
$\mu_{високий}(u_i)$	0,11	0,25	0,30	0,80	1

Таблиця 4.5

Функції належності терм-множини $H(t_2)$
(об'єктивність автора)

u_i , у. о.	1	2	3	4	5
$\mu_{низька}(u_i)$	1	0,85	0,45	0,35	0,10
$\mu_{середня}(u_i)$	0,15	0,44	1	0,85	0,15
$\mu_{висока}(u_i)$	0,06	0,12	0,50	0,60	1

Таблиця 4.6

Функції належності терм-множини $H(t_3)$
(інформативність контексту повідомлення)

u_i , у. о.	1	2	3	4	5
$\mu_{низька}(u_i)$	1	0,85	0,45	0,35	0,10
$\mu_{середня}(u_i)$	0,15	0,44	1	0,85	0,15
$\mu_{висока}(u_i)$	0,06	0,12	0,50	0,60	1

Таблиця 4.7

Функції належності терм-множини $H(l_1)$
(спростування і критика – частота)

u_i , у. о.	1	2	3	4	5
$\mu_{мала}(u_i)$	1	0,30	0,22	0,10	0,05
$\mu_{середня}(u_i)$	0,25	0,60	1	0,30	0,25
$\mu_{значна}(u_i)$	0,11	0,40	0,50	0,87	1

Таблиця 4.8

Функції належності терм-множини $H(l_2)$
(соціальна довіра)

u_i , у. о.	1	2	3	4	5
$\mu_{низька}(u_i)$	1	0,30	0,22	0,10	0,05
$\mu_{середня}(u_i)$	0,25	0,60	1	0,30	0,25
$\mu_{висока}(u_i)$	0,11	0,40	0,50	0,87	1

Наступний крок полягає у застосуванні таблиць 3.6-3.13 для вибору з певного інтервалу універсальної бази числових значень функцій належності усіх ЛЗ, означених відповідними нечіткими термами, для використання їх у ролі вихідних даних для розрахунку нечітких логічних рівнянь.

Для лінгвістичних змінних сформовано відповідну множину значень U :

$$\begin{aligned}
 &b_1 = 5; b_2 = 3; b_3 = 4; t_1 = 4; t_2 = 3; t_3 = 2; l_1 = 3; l_2 = 4. \\
 &\mu_{низька}(b_1) = 0,11; \mu_{середня}(b_1) = 0,12; \mu_{висока}(b_1) = 1; \\
 &\mu_{нечаста}(b_2) = 0,44; \mu_{часта}(b_2) = 1; \mu_{постійна}(b_2) = 0,68; \\
 &\mu_{мала}(b_3) = 0,11; \mu_{середня}(b_3) = 0,40; \mu_{велика}(b_3) = 0,80; \\
 &\mu_{низький}(t_1) = 0,30; \mu_{середній}(t_1) = 0,50; \mu_{високий}(t_1) = 0,80; \\
 &\mu_{низька}(t_2) = 0,45; \mu_{середня}(t_2) = 1; \mu_{висока}(t_2) = 0,50; \\
 &\mu_{низька}(t_3) = 0,85; \mu_{середня}(t_3) = 0,44; \mu_{висока}(t_3) = 0,12; \\
 &\mu_{мала}(l_1) = 0,22; \mu_{середня}(l_1) = 1; \mu_{значна}(l_1) = 0,50; \\
 &\mu_{низька}(l_2) = 0,10; \mu_{середня}(l_2) = 0,30; \mu_{висока}(l_2) = 0,87.
 \end{aligned}$$

Виконано обчислення функцій належності, що характеризують значення відповідних лінгвістичних змінних: B «організаційна складова ДТІП»; T «технічна складова ДТІП»; L «користувачська складова ДТІП». Для цього у нечіткі логічні рівняння, що стосуються вказаних ЛЗ, підставимо отримані вище значення функцій належності.

Отже, для лінгвістичної змінної B отримано такі функції належності:

$$\mu_{\text{низька}}(B) = 0,11 \wedge 0,44 \wedge 0,11 \vee 0,12 \wedge 0,44 \wedge 0,11 = 0,11;$$

$$\mu_{\text{середня}}(B) = 0,12 \wedge 1 \wedge 0,80 \vee 1 \wedge 1 \wedge 0,11 = 0,12;$$

$$\mu_{\text{висока}}(B) = 1 \wedge 0,56 \wedge 0,40 \vee 0,12 \wedge 0,56 \wedge 0,80 = 0,40.$$

Функції належності лінгвістичної змінної T містять такі значення:

$$\mu_{\text{низька}}(T) = 0,30 \wedge 0,50 \wedge 0,85 \vee 0,30 \wedge 0,45 \wedge 0,44 = 0,30;$$

$$\mu_{\text{середня}}(T) = 0,50 \wedge 1 \wedge 0,44 \vee 0,50 \wedge 0,50 \wedge 0,85 = 0,50;$$

$$\mu_{\text{висока}}(T) = 0,80 \wedge 1 \wedge 0,12 \vee 0,80 \wedge 0,50 \wedge 0,12 = 0,12.$$

Числові значення функцій належності лінгвістичної змінної L :

$$\mu_{\text{низька}}(L) = 0,22 \wedge 0,30 \vee 1 \wedge 0,10 = 0,22;$$

$$\mu_{\text{середня}}(L) = 0,50 \wedge 0,10 \vee 1 \wedge 0,30 = 0,30;$$

$$\mu_{\text{висока}}(L) = 0,50 \wedge 0,60 \vee 0,50 \wedge 0,15 = 0,50.$$

Виходячи з обчислених значень функцій належності лінгвістичних змінних, які характеризують окремі показники рівня достовірності текстових інформаційних повідомлень, що входять до класифікаційних груп змінних моделі логічного виведення, отримано функції належності, що описують лінгвістичну змінну найвищого рівня. Ця змінна відображає інтегральний показник рівня достовірності текстових інформаційних повідомлень з термами «низька», «середня», «висока».

$$\mu_{\text{низька}}(Q) = 0,11 \wedge 0,30 \wedge 0,22 \vee 0,11 \wedge 0,30 \wedge 0,22 = 0,11;$$

$$\mu_{\text{середня}}(Q) = 0,12 \wedge 0,30 \wedge 0,30 \vee 0,12 \wedge 0,50 \wedge 0,50 = 0,12;$$

$$\mu_{\text{висока}}(Q) = 0,40 \wedge 0,12 \wedge 0,30 \vee 0,40 \wedge 0,12 \wedge 0,50 = 0,12.$$

Визначення кількісного показника прогнозової достовірності текстових інформаційних повідомлень реалізовано за допомогою методу центру мас або центру ваги геометричної фігури, утвореної графіком функції належності, та відображає просторову інтерпретацію ступеня належності [50,84]. Цей метод є одним із найпоширеніших у теорії нечітких множин, оскільки забезпечує високу точність оцінювання шляхом врахування всіх можливих рівнів належності значень змінної.

Цей етап є зворотним до фазифікації та передбачає виконання дефазифікації нечіткої множини (4.1), використовуючи розраховані раніше значення функцій належності. Під час дефазифікації виконується перехід від нечіткого представлення інформації до конкретного числового значення, що дозволяє отримати об'єктивну оцінку рівня достовірності повідомлення. Дефазифікація, що базується на нечіткому логічному виведенні, дозволяє прогнозовано оцінити рівень достовірності повідомлень залежно від заданих вихідних параметрів. Завдяки такому підходу можна забезпечити гнучкість оцінювання та підвищити його адаптивність до різних типів текстових даних, що є важливим для інформаційних систем і процесів прийняття рішень.

Розрахунок прогнозного показника достовірності інформаційних повідомлень P здійснено за наведеною нижче формулою:

$$P = \frac{\sum_{i=1}^m \left[\underline{Q} + (i-1) \frac{\bar{Q} - \underline{Q}}{m-1} \right] \mu_i(Q)}{\sum_{i=1}^m \mu_i(Q)}, \quad (4.2)$$

де: $\underline{Q}, \bar{Q}$ – означають мінімальне і максимальне значення показника достовірності повідомлень; m – кількість нечітких термів лінгвістичної змінної Q . Для розрахунків задаємо такі вихідні величини: $m = 3$; $\mu_1(Q) = \mu_{\text{низька}}(Q)$, $\mu_2(Q) = \mu_{\text{середня}}(Q)$, $\mu_3(Q) = \mu_{\text{висока}}(Q)$. Нижню і верхню межі значень для

лінгвістичної змінної Q виразимо у відсотках: $\underline{Q}=1\%$; $\overline{Q}=100\%$.

Обчислення виконано за трьома точками встановленого інтервалу: $q_1=1$; $q_2=50$; $q_3=100$. Остаточно отримано таке значення прогнозного показника достовірності інформаційних повідомлень:

$$P = \frac{1 \cdot 0,11 + 50 \cdot 0,12 + 100 \cdot 0,12}{0,11 + 0,12 + 0,12} = 51,74\%.$$

Отриманий показник задіяний у пропоновану схему користувацького підходу до оцінювання повідомлень таким чином. Інтервал оцінювання ділимо на три частини, кожна з яких відповідає лінгвістичним оціночним термам рівня достовірності повідомлення: «низький», «середній», «високий». Рівень «низький» належить до інтервалу (1-40), «середній» – (41-80), «високий» – (81-100). Таким чином, отриманий вище результат вказує, що при заданих вхідних значеннях лінгвістичних змінних (факторів процесу оцінювання) розрахований показник свідчить про *середній* рівень достовірності інформаційного повідомлення.

Оскільки передбачуване значення показника рівня достовірності текстових інформаційних повідомлень отримано для конкретних параметрів вхідних даних (лінгвістичних змінних), можна стверджувати, що існує можливість керування процесом оцінювання повідомлень через введення наперед заданих комбінацій факторів достовірності. Виконання цього завдання потребує створення відповідного програмного забезпечення, яке враховуватиме послідовність і зміст основних етапів, пов'язаних із застосуванням методів нечіткої логіки.

Для цього необхідно розробити алгоритмічну базу, що забезпечить автоматизовану обробку впливових факторів достовірності із використанням математичних моделей та правил нечіткої логіки. Важливо також передбачити механізм адаптації системи до зміни параметрів вхідних даних, що дозволить підвищити ефективність її роботи. Крім того, інтеграція розробленого програмного рішення у вже наявні системи інформаційного аналізу забезпечить цілісність та системність у процесі оцінювання достовірності повідомлень.

4.2. Оцінювання рівнів достовірності текстових інформаційних повідомлень методом багатокритеріальної оптимізації

Для розв'язання озвученого завдання використано метод багатокритеріальної оптимізації, альтернативами в якому рівні достовірності текстових інформаційних повідомлень визначені лінгвістичними термами «низький», «середній», «високий». В даній ситуації для прийняття обґрунтованого рішення задано нечіткі відношення і попарні переваги часток ефективності факторів, віднесених до певного рівня достовірності повідомлень, заданих числом на відрізку $[0;1]$.

На основі теоретичних засад методу багатокритеріальної оптимізації опишемо суть алгоритму його експериментальної реалізації стосовно проблематики визначення рівня достовірності текстових повідомлень [46].

1. Вважатимемо рівні достовірності текстових інформаційних повідомлень множиною $X\{x_1, x_2, x_3\}$, елементи якої мають відповідне лінгвістичне трактування: x_1 – «низький»; x_2 – «середній»; x_3 – «високий». При цьому кожний з рівнів оцінюється за факторами такими нечіткими відношеннями: F_1 – сукупність переваг стосовно джерела інформації; F_2 – професіоналізму автора; F_3 – інформативності контексту. Відношенням F_j відповідатимуть вагові значення факторів w_j , $j=1,3$ та відповідні їм функції належності $\mu_j(x, y)$.

Знаходимо згортку відношень $Z_1 = \bigcap_{j=1}^3 F_j$.

2. В одержаній таким чином множині $\{X, Z_1\}$ визначаємо множину недомінованих рівнів достовірності Z_1^{nd} з такими функціями належності

$$\mu_{Z_1}^{nd}(x) = 1 - \sup_{y \in X} \left\{ \sum_{j=1}^4 \mu_{Z_1}(y, x) - \mu_{Z_1}(x, y) \right\}. \quad (4.3)$$

3. Вираз (4.3) стає підставою для адитивної згортки відношень Z_2 , функції належності якої обумовлюються такими виразами:

$$\mu_{Z_2}(x, y) = \sum_{j=1}^3 w_j \mu_j(x, y), \quad \sum_{j=1}^3 w_j = 1, \quad w_j \geq 0. \quad (4.4)$$

4. Для Z_2 формуємо множину недомінованих альтернатив, функції належності яких розраховуються за таким виразом:

$$\mu_{Z_2}^{nd}(x) = 1 - \sup_{y \in X} \left\{ \sum_{j=1}^4 \mu_{Z_2}(y, x) - \mu_{Z_2}(x, y) \right\}. \quad (4.5)$$

5. Розраховуємо значення функція належності спільної множини недомінованих альтернатив як перетин множин Z_1^{nd} та Z_2^{nd} , тобто $Z_{nd} = Z_1^{nd} \cap Z_2^{nd}$

$$\mu_{nd}(x) = \min \{ \mu_{Z_1}^{nd}(x), \mu_{Z_2}^{nd}(x) \}. \quad (4.6)$$

Результатом вважається лінгвістичне трактування рівня ДТІП, функція належності $\mu_{nd}(x)$ якого максимальна.

Враховуючи наведений алгоритм, для факторів достовірності повідомлень, що утворюють окрему множину Парето, задано альтернативні рівні достовірності, виражені множиною $X \{x_1, x_2, x_3\}$.

Недоміновані фактори множини Парето наведені у п 1 алгоритму. Додатково визначено відповідні нормалізовані вагові значення для згортки Z_2 : $w_1 = 0,5$; $w_2 = 0,3$; $w_3 = 0,2$.

Відношення переваг часток ефективності факторів у варіантах рівнів

достовірності повідомлень сформовано у вигляді таких комбінацій.

Джерело інформації (F_1): $x_1 < x_2$, $x_1 < x_3$, $x_2 < x_3$.

Професіоналізм автора (F_2): $x_1 < x_2$, $x_1 = x_3$, $x_2 > x_3$.

Інформативність контексту (F_3): $x_1 > x_2$, $x_1 > x_3$, $x_2 > x_3$.

Для означених факторів та прикріплених до них переваг побудовано матриці відношень [46,68,71]. Бінарну матрицю відношень фактора F_1 стосовно переваг джерела інформації в альтернативних варіантах рівнів достовірності текстових повідомлень задано табл. 4.9.

Таблиця 4.9

Матриця відношень фактора F_1 «Джерело інформації»

$\mu_{F_1}(x_i, x_j)$	x_i / x_j	x_1	x_2	x_3
	x_1	1	0	0
	x_2	1	1	0
	x_3	1	1	1

Таблиця 4.10

Матриця відношень фактора F_2 «Професіоналізм автора»

$\mu_{F_2}(x_i, x_j)$	x_i / x_j	x_1	x_2	x_3
	x_1	1	0	1
	x_2	0	1	1
	x_3	0	0	1

Таблиця 4.11

Матриця відношень фактора F_3 «Інформативність контенту»

$\mu_{F_3}(x_i, x_j)$	x_i / x_j	x_1	x_2	x_3
	x_1	1	1	1
	x_2	0	1	1
	x_3	0	0	1

Побудуємо згортку відношень $Z_1 = F_1 \cap F_2 \cap F_3$, матриця значень функції належності якої матиме наступну структуру.

Таблиця 4.12

	x_i / x_j	x_1	x_2	x_3
$\mu_{Z_1}(x_i, x_j)$	x_1	1	0	0
	x_2	0	1	0
	x_3	0	0	1

На основі табл. 4.4 та виразу (4.3) визначено множину недомінованих альтернатив.

$$\mu_{Z_1}^{n\partial}(x) = 1 - \sup_{y \in X} \{ \mu_{Z_1}(y, x) - \mu_{Z_1}(x, y) \}.$$

Для кожної з альтернатив отримано такі значення:

$$\mu_{Z_1}^{n\partial}(x_1) = 1 - \sup_{y \in X} \{ 0 - 0; 0 - 0 \} = 1;$$

$$\mu_{Z_1}^{n\partial}(x_2) = 1 - \sup_{y \in X} \{ 0 - 0; 0 - 0 \} = 1;$$

$$\mu_{Z_1}^{n\partial}(x_3) = 1 - \sup_{y \in X} \{ 0 - 0; 0 - 0 \} = 1.$$

За результатами обчислень $\mu_{Z_1}^{n\partial}(x) = [1; 1; 1]$.

Нечітке відношення переваги Z_2 визначено як адитивну згортку відношень [68,71] F_j , $j = 1, 3$ за формулою $Z_2 = \sum_{j=1}^3 w_j f_j(x)$. Виконано розрахунок функцій належності згортки Z_2 за формулою $\mu_{Z_2}(x_i, x_j) = \sum_{k=1}^3 w_k \mu_{F_k}(x_i, x_j)$. Отримані значення представлено у табл. 4.13.

Таблиця 4.13

Результуючі величини функцій належності згортки Z_2

	x_i / x_j	x_1	x_2	x_3
$\mu_{Z_2}(x_i, x_j)$	x_1	1	0.9	0.4
	x_2	0.3	1	0.5
	x_3	0.6	0.6	1

Обчислення згорток виконано на основі значень функцій належності, заданих

матрицями відношень у таблицях 4.9-4.13.

Для відношення Z_2 обчислено множину недомінованих альтернатив:

$$\mu_{Z_2}^{n\partial}(x_1) = 1 - \sup\{(0.9 - 0.3); (0.6 - 0.4)\} = 0.4;$$

$$\mu_{Z_2}^{n\partial}(x_2) = 1 - \sup\{(0.3 - 0.6); (0.5 - 0.6)\} = 1;$$

$$\mu_{Z_2}^{n\partial}(x_3) = 1 - \sup\{(0.9 - 0.5); (0.5 - 0.6)\} = 0.6.$$

Підсумком є отримання функції належності $\mu_{Z_2}^{n\partial}(x_i) = [0.4; 1; 0.6]$.

Завершальний етап представлено знаходженням згортки перетину недомінованих множин $Z_1^{n\partial}$ та $Z_2^{n\partial}$ тобто $Z_{n\partial} = Z_1^{n\partial} \cap Z_2^{n\partial}$, з функцією належності

$$\mu_Z^{n\partial}(x_i) = \min\{\mu_{Z_1}^{n\partial}(x_i), \mu_{Z_2}^{n\partial}(x_i)\}, \quad i = 1, 3. \quad (4.7)$$

З урахуванням того, що функція належності $\mu_{Z_1}^{n\partial}(x_i) = [1; 1; 1]$, отримано $\mu_Z^{n\partial}(x_i) = [0.4; 1; 0.6]$.

Функція належності згортки Z свідчить, що згідно із заданими вище відношеннями переваги корисності факторів в альтернативних варіантах та з урахуванням їх нормалізованих вагових значень апріорна достовірність текстового повідомлення відповідає альтернативі x_2 , тобто рівню «середній».

На завершення наведено інтерфейс програмного застосунку визначення прогностного рівня достовірності текстових інформаційних повідомлень за нечіткими відношеннями переваг мір ефективності факторів у визначених альтернативних варіантах.

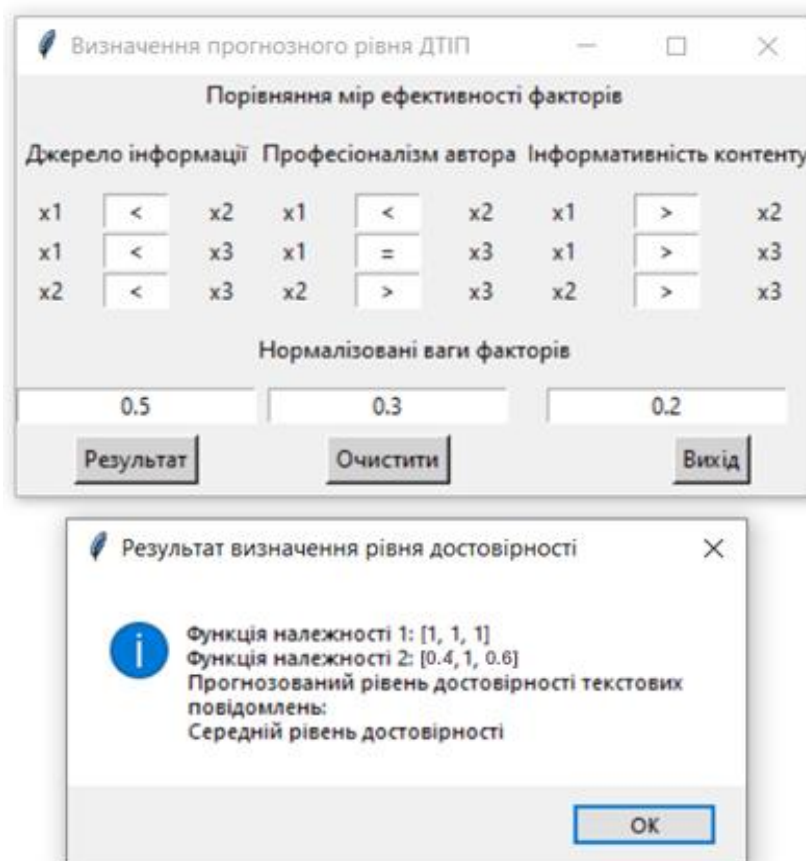


Рис. 4.1. Інтерфейс програмного застосунку визначення прогнозного рівня достовірності текстових інформаційних повідомлень

Представлений інтерфейс програмного застосунку реалізує етап визначення прогнозного рівня достовірності текстових інформаційних повідомлень на основі порівняння ефективності ключових факторів. Користувачеві надається можливість провести попарне порівняння таких чинників, як джерело інформації, професіоналізм автора та інформативність контенту. На основі введених даних розраховуються нормалізовані ваги факторів, які надалі застосовуються у процедурі нечіткого логічного виведення. Результатом обчислень є числове представлення функцій належності та лінгвістична оцінка рівня достовірності повідомлення, що виводиться у вигляді діалогового вікна. Програмний інтерфейс забезпечує базові функції обчислення, очищення введених даних і завершення роботи.

4.3. Функціональне моделювання процесу розроблення інформаційної технології оцінювання достовірності текстових повідомлень

З метою формалізації логіки функціонування інформаційної технології оцінювання достовірності текстових повідомлень доцільно застосувати методологію IDEF0, яка є ефективним засобом функціонального моделювання складних систем і процесів. Цей підхід забезпечує можливість ієрархічного подання процесу у вигляді послідовної сукупності взаємопов'язаних функцій, що сприяє більш глибокому розумінню структури та взаємодії окремих компонентів системи. У межах цієї методології побудовано контекстну діаграму рівня A0, яка відображає загальну мету та межі моделі, а також визначає вхідні й вихідні дані, механізми та керуючі впливи. Подальша деталізація реалізована шляхом декомпозиції на діаграму рівня A0, що конкретизує основні функціональні підсистеми, і відповідні діаграми нижчого рівня (A1, A2, A3, A4, A5), які послідовно розкривають логіку реалізації окремих етапів процесу створення інформаційної технології. Даний підхід дозволяє наочно і формалізовано представити ключові процедури, що забезпечують оцінювання достовірності текстових інформаційних повідомлень, і сприяє підвищенню прозорості, керованості та ефективності розроблюваної системи. [85,87].

Контекстна діаграма рівня A-0 відображає систему розроблення інформаційної технології оцінювання достовірності текстових повідомлень у взаємодії з зовнішнім середовищем. Відповідно до методології IDEF0, інформаційні потоки системи поділяються за типами: Вхід (Input), Керування (Control), Вихід (Output) та Механізм (Mechanism).

Контекстну діаграму представлено схематично на рис. 4.2. Вона містить умовні позначення, що відображають ключові елементи системи. Зокрема, I_1 — лінгвістичні характеристики, I_2 — джерело інформації (посилання), C_1 — нормативно-технічна та технологічна документація, C_2 — теорії, методи, методики, принципи, O_1 — оптимальна альтернатива реалізації, O_2 — інтегральний

показник рівня достовірності текстових повідомлень, M_1 — апаратне та програмне забезпечення, M_2 — дослідники, експерти у предметній галузі та інші зацікавлені учасники процесу [86,87].

Опишемо кожен граничну стрілку, враховуючи поділ за типами.

Граничні стрілки категорії «Вхід» (Input):

- I_1 (текстові повідомлення). Виступають базовим об'єктом обробки в межах інформаційної технології. Йдеться про різноманітні текстові матеріали, що потребують аналізу з метою перевірки їхньої достовірності. Ці повідомлення можуть містити як новини, так і інформаційні матеріали з відкритих джерел [87].

- I_2 (джерело інформації). Посилання на першоджерела або супровідну інформацію, яка слугує контекстом або підтвердженням для аналізованого текстового повідомлення. Цей вхід є критично важливим для оцінки достовірності, оскільки дозволяє здійснювати фактологічну перевірку [87].

Наведемо характеристику кожного з типів інформаційних потоків.

Граничні стрілки категорії «Контроль» (Control):

- C_1 (нормативно-технічна та технологічна документація). Регламентує процеси створення та функціонування інформаційної технології. Включає стандарти, методичні рекомендації, технічні умови та інші документи, що визначають порядок розроблення та експлуатації системи [87].

- C_2 (теорії, методи, методики, принципи). Формує концептуальну й методологічну основу для розроблення інформаційної технології. До цієї категорії належать сучасні підходи до оцінювання достовірності інформації. Так, наприклад, виокремлення та формалізація зв'язків між факторами визначення достовірності даних здійснюється за методами системного та матричного аналізу, теорією графів та семантичних мереж, логікою предикатів; створення моделі пріоритетного впливу факторів відбувається за теорією ієрархічних

багаторівневих систем. [87].

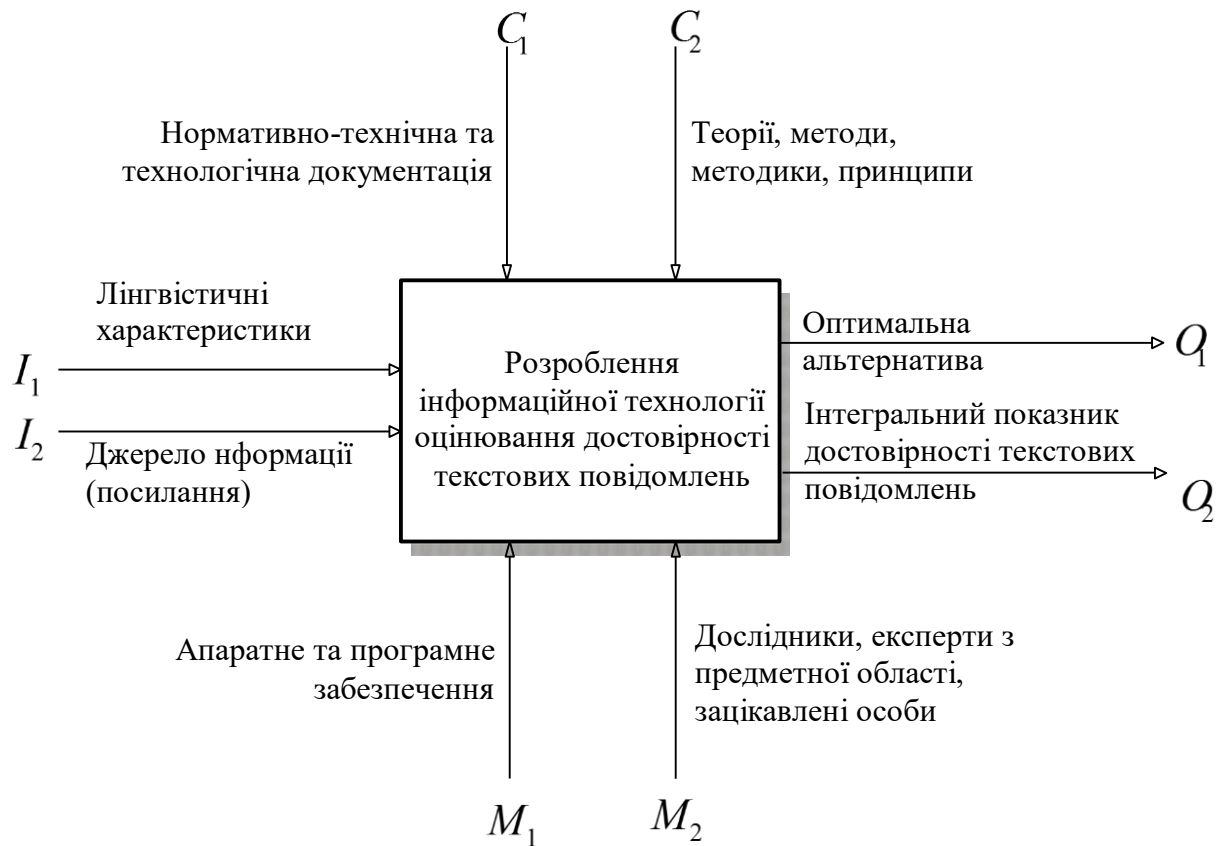


Рис. 4.2. Контекстна діаграма функціональної моделі процесу розроблення інформаційної технології оцінювання достовірності текстових повідомлень [87]

Граничні стрілки категорії «Вихід» (Output):

– O_1 (оптимальна альтернатива реалізації). Результатом роботи системи є обґрунтоване рішення щодо найбільш ефективного способу реалізації інформаційної технології оцінювання достовірності. Це може включати вибір технічної архітектури, алгоритмів, програмного забезпечення [86-88].

– O_2 (інтегральний показник достовірності текстових повідомлень). Комплексний кількісний або якісний показник, що характеризує достовірність аналізованих текстових повідомлень. Цей результат дозволяє користувачам або системам ухвалювати рішення щодо довіри до отриманої інформації. Визначення

кількісного показника якості проводиться із застосуванням методів і інструментів нечіткої логіки [86, 87, 89, 90].

Граничні стрілки категорії «Механізми» (Mechanism):

– M_1 (апаратне та програмне забезпечення). Процеси пошуку, обробки, збереження та передачі інформації реалізуються з використанням сучасних технічних і програмних ресурсів [86, 87].

– M_2 (дослідники, експерти з предметної області, зацікавлені особи). Ключові учасники процесу розроблення та експлуатації системи. Сюди відносяться фахівці, що володіють експертними знаннями у сфері достовірності інформації, лінгвістики, інформаційних технологій, а також кінцеві користувачі та стейкхолдери [86, 87].

В ході деталізації контекстної діаграми розроблено діаграму рівня A0, яка відображає структуровану послідовність основних функціональних етапів процесу оцінювання достовірності текстових інформаційних повідомлень.

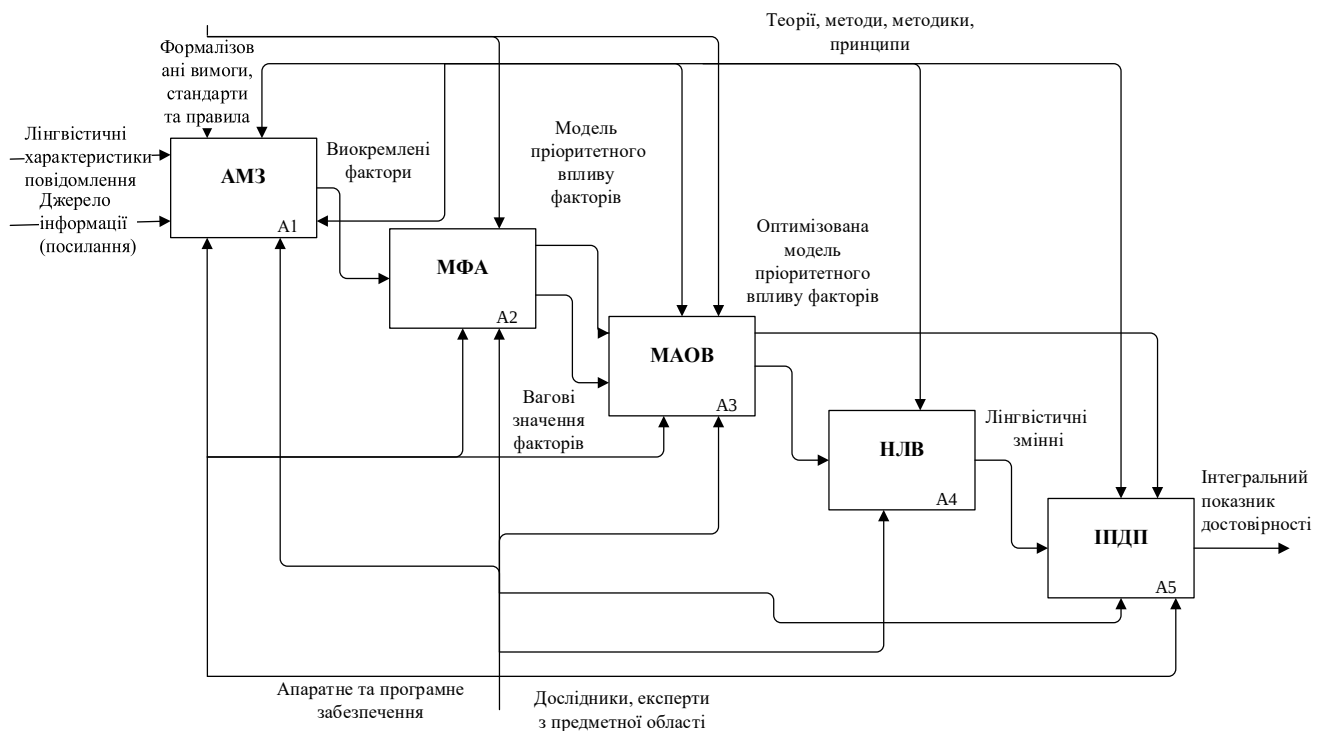


Рис. 4.3. Функціональна модель процесу оцінювання достовірності текстових інформаційних повідомлень

У структурі діаграми виокремлено такі функціональні блоки:

– АМЗ (аналіз методів і засобів). На цьому етапі здійснюється комплексне дослідження теоретико-методичних засад процесу оцінювання достовірності повідомлень. Особливу увагу приділено ідентифікації структурних компонентів процесу та формалізації його послідовності шляхом побудови відповідних функціональних моделей [86, 87].

– МФА (методи факторного аналізу). Передбачає застосування факторного підходу для формалізації взаємозв'язків між інформаційними та контекстуальними ознаками повідомлень. Побудова семантичної мережі здійснюється із залученням елементів логіки предикатів, що дозволяє відобразити структуру залежностей між факторами. Додатково реалізується процедура визначення рівнів домінантності окремих факторів із використанням методів ієрархічного математичного моделювання та уточнення результатів за допомогою методів ранжування [86, 87].

– МАОВ (моделювання альтернатив оптимальних варіантів). Даний етап спрямований на побудову оптимізованої моделі оцінювання достовірності, що передбачає визначення вагових коефіцієнтів факторів із урахуванням їх впливу на кінцевий результат. Застосування відповідного математичного апарату дозволяє усунути неоднозначність у встановленні пріоритетності факторів, зокрема уникнути ситуацій із однаковими рівнями значущості [87].

– НЛВ (нечітке логічне виведення). Передбачає розроблення нових та використання існуючих методів нечіткої логіки для моделювання альтернативних сценаріїв процесу оцінювання достовірності. Здійснюється порівняльний аналіз результатів, отриманих за різними підходами, що дозволяє забезпечити валідацію моделі та підтвердити її адекватність у контексті практичних застосувань [86-88].

– ПДП (інтегральний показник достовірності повідомлень). На основі розроблених методів, побудованих моделей і результатів аналізу здійснюється числова оцінка достовірності текстових інформаційних повідомлень. Визначення

інтегрального показника виконується з використанням методів нечіткої логіки та математичних моделей прогнозування, що забезпечує комплексну кількісну характеристику рівня достовірності з урахуванням множини релевантних факторів [87, 89, 90].

Запропонована структура функціональних блоків діаграми A0 забезпечує системний підхід до оцінювання достовірності текстових інформаційних повідомлень із використанням методів факторного аналізу, математичного моделювання та нечіткої логіки [91,92].

4.4. Структурно-функціональна модель інформаційної технології оцінювання достовірності текстових повідомлень

Узагальнення результатів попередніх розділів дисертаційної роботи зумовлює необхідність побудови структури інформаційної технології для оцінки рівня достовірності текстових повідомлень. Її формування потребує цілісного концептуального узагальнення одержаних даних, встановлення логічних зв'язків між етапами дослідження, визначення їх змістового наповнення, а також окреслення елементів теоретичної та прикладної новизни. Особливої уваги набуває експериментальна інтерпретація отриманих результатів, що забезпечує внутрішню узгодженість моделі та її відповідність поставленим цілям. Визначальним є те, що інформаційна технологія розглядається як інтегрована система, яка використовує інструменти прогнозування потенційної достовірності повідомлень до їх появи в медійному просторі. Такий підхід забезпечує єдність прогностичного та діагностичного підходів, що дозволяє підвищити обґрунтованість оцінювання, зменшити вплив суб'єктивних чинників та забезпечити вищу точність аналізу. Інформаційна технологія, сконструйована у такий спосіб, набуває функціональної завершеності, що дозволяє розглядати її як універсальний інструмент, придатний до адаптації під широкий спектр завдань.

Виокремимо дослідницькі та експериментальні етапи процесу оцінювання рівня достовірності текстових інформаційних повідомлень.

Етап 1. Аналіз методів та засобів процесу оцінювання достовірності текстових інформаційних повідомлень.

Е 1.1. Аналіз сучасного стану проблематики оцінювання рівня достовірності інформаційних повідомлень, що містить огляд систем розпізнавання достовірності текстових інформаційних повідомлень.

Е 1.2. Обґрунтування доцільності застосування засобів факторного аналізу для дослідження проблематики оцінювання достовірності текстових інформаційних повідомлень.

Е 1.3. Аналіз процесів контролю та управління ходом встановлення рівня достовірності отриманих даних. Експертне виокремлення факторів процесів.

Е 1.4. Опис характеристик особливостей утворення та використання лінгвістичних змінних у процесі оцінювання рівня достовірності даних.

Е 1.5. Побудова моделі виконання дисертаційного дослідження, орієнтованої на розроблення методів, моделей та інструментів для формування та оцінювання показника ДТІП із використанням теорії нечітких множин.

Етап 2. Розроблення методів факторного аналізу.

Е 2.1. Експертне виокремлення факторів та просторове відображення зв'язків між ними за допомогою семантичної мережі. Формалізований опис мережі засобами мови предикатів.

Е 2.2. Синтез багаторівневих моделей факторів впливу на достовірність повідомлень за методами математичного моделювання ієрархій та ранжування.

Е 2.3. Оптимізація моделі впливових факторів на процес формування показника ДТІП з використанням методу та матриці попарних порівнянь.

Етап 3. Моделювання альтернативних та оптимального варіантів процесу формування показника достовірності інформаційних повідомлень.

Е 3.1. Розроблення теоретичних засад альтернативних варіантів процесу формування показника ДТІП на основі факторів множини Парето.

Е 3.2. Розрахунок для факторів множини Парето функцій корисності альтернативних варіантів процесу методом лінійного згортання критеріїв.

Е 3.3. Визначення оптимального варіанту реалізації процесу встановлення достовірності інформаційних повідомлень за максимальним значенням об'єднаного функціоналу

Етап 4. Застосування нечіткої логіки для прогностичного оцінювання рівня достовірності текстових інформаційних повідомлень.

Е 4.1. Структурування процесу формування показника рівня ДТІП через створення основних складових: організаційної, технічної та користувацької.

Е 4.2. Формування бази даних, що обумовлює зв'язок між лінгвістичною природою ЛЗ, значеннями терм-множини та лінгвістичними термами.

Е 4.3. Розроблення методу нечіткого логічного виведення з урахуванням груп лінгвістичних змінних.

Е 4.4. Розрахунок значень та графічне відображення функцій належності лінгвістичних змінних з використанням гаусівських функцій.

Е 4.5. Проектування нечітких матриць знань, формування нечітких логічних рівняння для лінгвістичних змінних моделі логічного виведення.

Е 4.6. Обчислення показника прогностичного рівня ДТІП на основі значень функцій належності параметрів лінгвістичних змінних.

Етап 5. Реалізація інформаційної технології оцінювання рівня ДТІП

Е 5.1. Розрахунок показника прогностичного рівня достовірності текстових інформаційних повідомлень.

Е 5.2. Розроблення функціональної моделі оцінювання достовірності інформаційних повідомлень

Е 5.3. Експериментальне дослідження алгоритмічного і програмного забезпечення процесу.

Узагальнену структурно-функціональну модель інформаційної технології оцінювання достовірності текстових інформаційних технологій, з урахуванням виокремлених функціональні блоків рис 4.3, наведено на рис. 4.4.



Рис. 4.4. Структурно-функціональна модель інформаційної технології прогнозного оцінювання достовірності текстових повідомлень

Отримані результати сформульовано у вигляді таких тверджень.

1. Структурно-функціональна модель забезпечила інтегрування сучасних методів аналізу текстових даних для оцінювання їх достовірності. Модель обумовлює комплексний підхід до аналізу повідомлень за допомогою засобів нечіткої логіки та імітаційного моделювання.

2. Ідентифіковано ключові фактори впливу на рівень достовірності текстової інформації, такі як семантична узгодженість, інформативність та стильові характеристики тексту, джерело походження даних та їх контекстна релевантність, професіоналізм і об'єктивність автора, спростування і критика та соціальна довіра, що забезпечує автоматизоване оцінювання з високим рівнем точності.

3. Сформовано підходи до інтеграції в розроблену систему, що дозволяє ефективно застосовувати технологію в різних сферах, таких як видавництва різного рівня, освіта, маркетинг та державне управління.

4. Запропонована технологія має високий практичний потенціал, оскільки сприяє підвищенню рівня інформаційної безпеки, протидії дезінформації та зміцненню довіри до джерел інформації.

6. Ефективність методів оцінено на прикладі розрахунку показника прогнозованого рівня достовірності текстових інформаційних повідомлень для експериментальних значень функцій належності лінгвістичних змінних досліджуваного процесу. Отримані результати підтверджують здатність запропонованої технології адаптуватися до різноманітних факторів впливу на достовірність текстової інформації та забезпечувати точність прогнозування достовірності даних.

7. Модель та її прикладна компонента – запроєктоване програмне рішення – забезпечують оцінювання рівня достовірності текстових інформаційних повідомлень оптимізують процеси перевірки інформації в умовах великого обсягу даних, що є важливим аспектом у сучасному видавничому середовищі.

8. Практичні результати дослідження можуть бути використані у розробленні програмного забезпечення для медіа корпорацій, видавництв, освітніх проєктів та систем у сфері безпеки.

Таким чином, розроблена інформаційна технологія є надійним інструментом для вирішення актуальних задач в умовах зростаючого обсягу інформаційних потоків, надаючи користувачам можливість приймати обґрунтовані рішення, базуючись на достовірній інформації.

4.5. Порівняння методів оцінювання достовірності текстових повідомлень

Ефективність будь-якої інформаційної технології значною мірою визначається її здатністю демонструвати конкурентні переваги порівняно з наявними аналогами. У контексті оцінювання достовірності текстових повідомлень це означає не лише досягнення високої точності розпізнавання фейкових даних, а й забезпечення адаптивності до контексту та здатності працювати з неоднозначними даними [93,94]. Нами здійснено порівняльний аналіз різних підходів до визначення рівня достовірності текстових повідомлень, зокрема традиційних статистичних моделей, методів машинного навчання та запропонованої в дослідженні гібридної технології прогнозного оцінювання, що поєднує багатofакторний аналіз із засобами нечіткої логіки.

Результати аналізу та порівняння (див таблицю 4.14) засвідчують конкурентоспроможність розробленого підходу, що проявляється у підвищеній точності, стабільності результатів та здатності враховувати впливові фактори. Інтеграція методів нечіткого моделювання дозволила ефективно формалізувати експертні оцінки, зменшити рівень невизначеності та забезпечити більш гнучке реагування на складні інформаційні умови [95,96]. На відміну від методів, орієнтованих винятково на аналіз статистичних закономірностей, розроблена технологія орієнтується на прогнозу оцінку майбутніх повідомлень до їх поширення в медійному просторі.

Таблиця 4.14

Порівняльні характеристики методів оцінювання
достовірності текстових повідомлень

Критерій порівняння	Розроблена інформаційна технологія (методи багатофакторного аналізу та нечіткої логіки)	Статистичні методи (класифікація регресія)	Методи машинного навчання (SVM, Decision Tree, Random Forest)	Нейромережові методи (BERT, GPT, LSTM)
Комплексність аналізу	Висока — врахування кількох факторів одночасно, зокрема джерела, стилістики, структури, експертних оцінок	Обмежена — аналіз окремих ознак або показників	Середня — врахування множини ознак за рахунок побудови моделі	Висока — комплексний аналіз контексту та семантики тексту
Обробка невизначеності	Так — використання нечіткої логіки для врахування неоднозначності інформації	Ні — базуються на чітких числових значеннях	Частково — за рахунок ймовірнісних підходів	Частково — внутрішні механізми можуть враховувати невизначеність
Зрозумілість результатів (Explainability)	Висока — результати легко інтерпретуються завдяки формальним правилам і нечітким множинам	Висока — прості моделі легко пояснити	Середня — залежить від складності моделі	Низька — результати «чорного ящика» складно пояснити
Гнучкість налаштування	Висока — адаптація набору факторів і правил під різні задачі	Середня — налаштування параметрів моделей	Середня — налаштування гіперпараметрів	Обмежена — залежність від архітектури та обсягу даних

Вимоги до обсягу даних	Помірні — технологія може працювати навіть із обмеженим набором даних	Низькі — можливе застосування за малого обсягу даних	Середні — потребують достатнього обсягу даних для навчання	Високі — необхідні великі корпуси текстів для навчання
Стійкість до маніпуляцій та фейків	Висока — завдяки багатофакторному підходу складніше обійти систему	Низька — легко вводити в оману через штучні ознаки	Середня — залежить від якості навчання	Середня — можлива вразливість до замаскованих фейків
Сфера застосування	ЗМІ, видавничі системи, системи інформаційної безпеки	Академічні дослідження, базовий аналіз	Комерційні застосунки, автоматизовані системи	Великомасштабні платформи, автоматизація перевірки фактів

Таким чином, обґрунтовано доцільність застосування запропонованого методу як надійного засобу для верифікації текстових повідомлень, що дозволяє суттєво підвищити рівень інформаційної безпеки та мінімізувати ризики деструктивного впливу недостовірної інформації на користувача.

Розроблена інформаційна технологія демонструє баланс між пояснюваністю, гнучкістю та стійкістю до фейкових повідомлень, що вигідно вирізняє її серед традиційних статистичних чи машинних методів. На відміну від нейромережевих моделей, вона не потребує надвеликих обсягів даних для ефективного функціонування, зберігаючи при цьому можливість враховувати багатовимірні, якісні та нечіткі фактори, притаманні реальним комунікаційним видавничим середовищам.

Висновки до розділу 4

1. Реалізовано процес розрахунку значення показника прогнозного рівня достовірності текстових інформаційних повідомлень на основі обчислених функцій належності, що базуються на попередньо встановлених значеннях

лінгвістичних змінних.

2. Здійснено дефазифікацію нечіткої множини, реалізовану на основі нечіткого логічного виведення і вихідних даних лінгвістичних змінних:

$$b_1 = 5 \text{ у.о.}; b_2 = 3 \text{ у.о.}; b_3 = 4 \text{ у.о.}; t_1 = 4 \text{ у.о.}; t_2 = 3 \text{ у.о.}; t_3 = 2 \text{ у.о.}; l_1 = 3 \text{ у.о.}; l_2 = 4 \text{ у.о.}$$

Для наведеного варіанту вихідних даних за методом центра мас розраховано прогнозоване значення показника рівня достовірності текстових інформаційних повідомлень, величина якого $P = 51,74\%$, що свідчить про середній рівень достовірності.

3. Виконано імітаційне моделювання процесу оцінювання достовірності текстових інформаційних повідомлень із використанням методу багатокритеріальної оптимізації. Для його реалізації використано множину впливових Парето-факторів, до яких віднесено «джерело інформації», «професіоналізм автора», «інформативність контенту».

4. Розроблено програмний застосунок, що автоматизує процедуру прийняття рішення щодо рівня достовірності повідомлень на основі вхідних даних, отриманих способом логічного вираження попарних переваг між альтернативними рівнями достовірності текстових повідомлень. Наведено інтерфейс користувача, що забезпечує введення формування вхідних даних та візуалізацію результатів.

5. Здійснено функціональне моделювання процесу створення інформаційної технології, реалізоване з використанням методології IDEF0. Це дозволяє відобразити логіку побудови технології у вигляді структурних взаємозалежних процесів, уточнити рольових учасників, інформаційні потоки, вхідні й вихідні параметри, а також ресурси, необхідні для реалізації кожного з етапів.

6. Розроблено інформаційну технологію оцінювання достовірності текстових повідомлень, структурно-функціональна модель якої містить такі етапи: аналіз методів та засобів процесу оцінювання ДТІП; розроблення методів факторного аналізу; моделювання альтернативних та оптимального варіантів

процесу формування показника достовірності інформаційних повідомлень; застосування нечіткої логіки для прогностного оцінювання рівня достовірності текстових інформаційних повідомлень; реалізація інформаційної технології оцінювання рівня ДТП.

7. Виконано процедуру порівняльного оцінювання методів достовірності текстових інформаційних повідомлень, до яких віднесено такі засоби: розроблена інформаційна технологія (метод багатофакторного аналізу та нечіткої логіки); статистичні методи (класифікація, регресія); методи машинного навчання (SVM, Decision Tree, Random Forest); Нейромережеві методи (BERT, GPT, LSTM).

ВИСНОВКИ

У дисертаційній роботі розв'язано актуальне науково-прикладне завдання прогнозного оцінювання достовірності текстових інформаційних повідомлень на основі факторного аналізу їх вірогідності із застосуванням інформаційної технології, розробленої на засадах теорії моделювання, дослідження операцій та апарату нечітких множин.

У процесі дослідження отримано такі основні результати.

1. Виконано аналіз методів, засобів та технологій оцінювання рівня достовірності текстових інформаційних повідомлень. Відзначено, що вивчення і дослідження вказаної тематики вимагає комплексного підходу для розв'язання складних проблем, пов'язаних з моніторингом інформації в сучасному інформаційному середовищі

2. Обґрунтовано доцільність застосування наведених засобів факторного аналізу для дослідження проблематики оцінювання достовірності текстових інформаційних повідомлень.

3. Розроблено багаторівневу модель виконання дисертаційного дослідження, орієнтовану на поєднання двох підходів, реалізованих створеними інформаційною та програмною компонентами, здатними ефективно визначати рівні достовірності текстових повідомлень у прогнозованому варіанті.

4. Розроблено графічно орієнтовану модель у вигляді семантичної мережі, що уможливила формалізоване відображення структурних та лінгвістичних відношень між факторами достовірності текстових повідомлень. Виконано опис семантичної мережі з використанням елементів логіки предикатів, що забезпечує можливість застосування математичного апарату теорії моделювання для вивчення особливостей впливу вказаних чинників на процес формування показника достовірності текстових повідомлень.

5. Розроблено на основі використання методів ранжування та визначення вагомості предикатів оптимізовану багаторівневу модель факторів достовірності текстових повідомлень, поетапний процес створення якої передбачає отримання

вагових значень факторів за критеріями: максимального власного значення $\lambda_{\max} = 8,39$ додатної обернено-симетричної матриці; індексу узгодженості $IU = 0,056$; відношення узгодженості $WU = 0,04$.

6. Запроектовано на підставі методу лінійного згортання критеріїв альтернативні варіанти визначення рівня довіри до інформації з використанням функцій корисності, часток ефективності факторів множини Парето в альтернативах та часткових функціоналів. Розраховано оптимальний варіант процесу формування показника достовірності текстових повідомлень за критерієм максимального значення цільового функціоналу, величина якого $U_1 = 0.480$.

7. Удосконалено метод формування показника достовірності текстових інформаційних повідомлень та визначення їх вірогідності, у якому реалізовано фазифікацію як процес переходу від числових даних до нечітких оцінок на основі лінгвістичних змінних, а також дефазифікацію як зворотну процедуру перетворення нечітких результатів у конкретні числові значення.

8. Розроблено метод логічного виведення, багаторівнева модель якого відображає логіку процесу формування показника прогностного рівня достовірності текстових повідомлень з урахуванням класифікаційних груп лінгвістичних змінних. Запроектовано нечіткі матриці знань і нечіткі логічні рівняння.

9. Здійснено дефазифікацію нечіткої множини, реалізовану на основі нечіткого логічного виведення і вихідних даних лінгвістичних змінних:

$$b_1 = 5 \text{ у.о.}; b_2 = 3 \text{ у.о.}; b_3 = 4 \text{ у.о.}; t_1 = 4 \text{ у.о.}; t_2 = 3 \text{ у.о.}; t_3 = 2 \text{ у.о.}; l_1 = 3 \text{ у.о.}; l_2 = 4 \text{ у.о.}$$

Для наведеного варіанту вхідних даних за методом центра мас розраховано прогнозоване значення показника рівня достовірності текстових інформаційних повідомлень, величина якого $P = 51,74\%$, що свідчить про середній рівень достовірності.

10. Розроблено програмний застосунок автоматизації процесу оцінювання

достовірності текстових інформаційних повідомлень із використанням методу багатокритеріальної оптимізації. Для його реалізації використано множину впливових Парето-факторів, до яких віднесено «джерело інформації», «професіоналізм автора», «інформативність контенту».

11. Здійснено функціональне моделювання процесу створення інформаційної технології, реалізоване з використанням методології IDEF0. Це дозволяє відобразити логіку побудови технології у вигляді структурних взаємозалежних процесів, уточнити рольових учасників, інформаційні потоки, вхідні й вихідні параметри, а також ресурси, необхідні для реалізації кожного з етапів.

12. Розроблено інформаційну технологію оцінювання достовірності текстових повідомлень, структурно-функціональна модель якої містить такі етапи: аналіз методів та засобів процесу оцінювання ДТІП; розроблення методів факторного аналізу; моделювання альтернативних та оптимального варіантів процесу формування показника достовірності інформаційних повідомлень; застосування нечіткої логіки для прогностного оцінювання рівня достовірності текстових інформаційних повідомлень; реалізація інформаційної технології оцінювання рівня достовірності текстових інформаційних повідомлень.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Гарасимчук О., Наконечний Ю., Луковський Т., Андріїв Р., Наконечний Т. Захищене зберігання даних із використанням технології блокчейн ETHEREUM. *Ukrainian Information Security Research Journal*. 2024. Том 26 № 1. С. 213-220. URL: <https://doi.org/10.18372/2410-7840.26.18845> (дата звернення: 07.02.2025).
2. Krishnamurthi, Dr. Ilango, Dr. Santhi V, and Madhumitha N H. Empirical Analysis for Classification of Fake News through Text Representation. *Journal of Information Technology and Digital World*. 2024. Vol. 6, no. 1. P.27-45 URL: <https://doi.org/10.36548/jitdw.2024.1.003> (дата звернення: 07.02.2025).
3. Ткаченко О., Ільєнко А., Улічев О., Мелешко Є., Смірнов О. Правові засади поширення інформаційних впливів в соціальних мережах. *Електронне фахове наукове видання Серія «Кибербезпека: освіта, наука, техніка»*, 2024. №2(26), С.170–188. URL: <https://doi.org/10.28925/2663-4023.2024.26.685> (дата звернення: 07.02.2025).
4. Mishra S., Shukla P., Agarwal R. Analyzing machine learning enabled fake news detection techniques for diversified datasets. *Wireless Communications and Mobile Computing*, 2022 (1) 1575365. URL: <https://doi.org/10.1155/2022/1575365>. (дата звернення: 07.02.2025).
5. Dixit D. K., Bhagat A., Dangi D. Automating fake news detection using PPCA and levy flight-based LSTM. *Soft Comput.*, 2022. Vol. 26, Issue 22, P. 12545–12557. URL: <https://doi.org/10.1007/s00500-022-07215-4>. (дата звернення: 08.02.2025)
6. Wang Y., Wang L., Yang Y., Zhang Y. Detecting fake news by enhanced text representation with multi-EDU-structure awareness. *Expert Syst. 2022. Appl.*, 206 p. URL: <https://doi.org/10.1016/j.eswa.2022.117781>. (дата звернення: 17.02.2025).
7. Dimitrios M., Kanakaris N., Varlamis I. Detection of fake news campaigns using graph convolutional networks. *Int. J. Inf. Manag. Data Insights*, 2022. 100104. Vol. 2, no. 2. URL: <https://doi.org/10.1016/j.jjime.2022.100104>. (дата звернення: 17.02.2025).

8. Rastogi S., Bansal D. Disinformation detection on social media: An integrated approach. *Multimedia Tools and Applications*, 2022. Vol. 81, Issue 28, P. 40675–40707. URL: <https://doi.org/10.1007/s11042-022-13129-y>. (дата звернення: 27.02.2025).

9. Rath B., Salecha A., Srivastava J. Fake news spreader detection using trust-based strategies in social networks with bot filtration. *Soc. Netw. Anal. Min.*, 2022. Vol. 12, Issue 1. URL: <https://doi.org/10.1007/s13278-022-00890-z> (дата звернення: 19.02.2025).

10. Sousa, S., Bates, N. Factors influencing content credibility in Facebook's news feed. *Human-Intelligent Systems Integration*. 2021. Vol. 3, P. 69–78. URL: <https://doi.org/10.1007/s42454-021-00029-z> (дата звернення: 21.02.2025)

11. Chen, L., Chen J., Xia C. Social network behavior and public opinion manipulation. *Journal of Information Security and Applications*. 2022. Vol. 64. 103060. URL: <https://doi.org/10.1016/j.jisa.2021.103060>.
(дата звернення: 21.02.2025).

12. Basol, M., Roozenbeek, J., Sander van der Linden. Good news about bad news: Gamified inoculation boosts confidence and cognitive immunity against fake news. *Journal of Cognition*, 2020. Jan 10; 3(1):2. URL: <https://doi.org/10.5334/joc.91>
(дата звернення: 22.02.2025).

13. Rampersad G., Althiyabi T. Fake news: Acceptance by demographics and culture on social media. *Journal of Information Technology & Politics* 2020. Vol. 17. no 1. P. 1-11. URL: <https://doi.org/10.1080/19331681.2019.1686676>
(дата звернення: 24.02.2025).

14. Jones-Jang, S. M., Mortensen, T., & Liu, J. Does media literacy help identification of fake news? Information literacy helps, but other literacies don't. *American Behavioral Scientist*, 2021. Vol. 65, no. 2, P. 371–388.
URL: <https://doi.org/10.1177/0002764219869406>. (дата звернення: 24.02.2025).

15. Ozkan M., Akin E., Aslan O. et al. A comprehensive survey: Evaluating the

efficiency of artificial intelligence and machine learning techniques on cyber security solutions. In *IEEE Access*, 2024. Vol. 12. P. 12229-12256. DOI: 10.1109/ACCESS.2024.3355547. (дата звернення: 25.02.2025).

16. Silva M. J. Digital media and misinformation: An outlook on multidisciplinary strategies against manipulation. *J Comput Soc Sc* 5, 2022. P. 123–159. URL: <https://doi.org/10.1007/s42001-021-00118-8> (дата звернення: 25.02.2025).

17. Ziad A., Mohammed I., Al-Saleh, Alawneh L., Jararweh J., Gupta B. Live forensics of software attacks on cyber–physical systems. *Future Generation Computer Systems*. Vol. 108, July 2020. P. 1217-1229
URL: <https://doi.org/10.1016/j.future.2018.07.028> (дата звернення: 27.02.2025).

18. Asghar M. Z., Ullah A., Ahmad S., Khan A., Opinion spam detection framework using hybrid classification scheme, *Soft Computing*. 2020. Vol. 24, No. 5, P. 3475–3498. URL: <https://doi.org/10.1007/s00500-019-04107-y> (дата звернення: 28.02.2025).

19. Pummy D., Kaur A., Iwendi C., Mohan S. A scientometric analysis of deep learning approaches for detecting fake news. *Electronics* 2023. 12(4). 948.

URL: <https://doi.org/10.3390/electronics12040948> (дата звернення: 28.02.2025).

20. Khanam, Zeba, et al. Fake news detection using machine learning approaches. *IOP conference series: materials science and engineering*. 2021. Vol. 1099. No. 1. IOP Publishing. URL: <https://doi.org/10.1088/1757-899x/1099/1/012040> (дата звернення: 28.02.2025).

21. Ilnytska U. Information security of Ukraine: modern challenges, threats and countermeasures against negative information and psychological influences. *Political Sciences*, 2016. No. 1(2), P. 27-32.

22. Brzhevskaya Z., Haydur H., Anosov A. Impact on the credibility of information as a threat to the information space. *Cyber security: education, science, technology*, 2018. No. 2(2), P. 105-112. URL: <https://doi.org/10.28925/2663-4023.2018.2.105112>

(дата звернення: 08.03.2025).

23. Korobchynskyi M., Rudenko M., Dereko V., Kovtun O., Zaitsev O. Optimization of Data Preprocessing Procedure in the Systems of High Dimensional Data Clustering. *Lecture Notes on Data Engineering and Communications Technologies*. 2023. Vol. 149. Pp. 449–461.

24. Brzhevskaya Z., Haydur H., Dovzhenko N., Anosov A. Criteria for monitoring the reliability of information in the information space. *Cyber security: education, science, technology*, 2019. No. 1(5), P.105-112. URL: <https://doi.org/10.28925/2663-4023.5.52.60>. (дата звернення: 10.03.2025).

25. Marcheniuk, M. S., Kozachok, V. A., Bogdanov, O. M., Brzhevskaya, Z. M. Analysis of methods of detecting disinformation in social networks using machine learning. *Cyber security: education, science, technology*, 2023. No. 2(22), P. 148-155. URL: <https://doi.org/10.28925/2663-4023.2023.22.148155>. (дата звернення: 10.03.2025).

26. Піх І. В., Сеньківський В. М., Андрійів Р. Р. Оптимізована модель факторів достовірності текстових даних. *Науковий вісник Національного лісотехнічного університету України*. 2024, т. 34. 2024. С. 78-85.

URL: <https://doi.org/10.36930/40340410> (дата звернення: 12.03.2025).

27. Thorne, J., Vlachos, A. *Automated Fact Checking: Task Formulations, Methods and Future Directions*. *Proceedings of the 27th International Joint Conference on Artificial Intelligence (IJCAI)*. 2018. P. 4510-4516. URL: <https://www.ijcai.org/proceedings/2018/627> (дата звернення: 14.03.2025).

28. Conroy, N. J., Rubin, V. L., & Chen, Y. Automatic deception detection: Methods for finding fake news. *Proceedings of the Association for Information Science and Technology*, 2015. 52(1), P. 1–4.

URL: <https://doi.org/10.1002/pa2.2015.145052010082> (дата звернення: 14.03.2025).

29. Wu, L., Morstatter, F., Carley, K. M., & Liu, H. *Misinformation in Social Media: Definition, Manipulation, and Detection*. *ACM SIGKDD Explorations*

Newsletter, 2019. 21(2), P. 80–90. URL: <https://doi.org/10.1145/3373464.3373475> (дата звернення: 22.03.2025).

30. Корж О., Коровай В. Фейковий контент: види, ознаки, шляхи виявлення. *Освіта. Інноватика. Практика*, 2023 11(7), С. 37–42 <https://doi.org/10.31110/2616-650X-vol11i7-005>

31. Zhou, X., Zafarani, R. Fake News: A Survey of Fake News: Fundamental Theories, Detection Methods, and Opportunities. *ACM Computing Surveys (CSUR)*, 2020. Vol. 53, Issue 5 Article No.: 109, P 1 – 4. URL <https://doi.org/10.1145/3395046> (дата звернення: 23.03.2025).

32. Baly, R., Karadzhov, G., Alexandrov, D., Glass, J., & Nakov, P. What Was Written vs. Who Read It: News Media Profiling Using Text Analysis and Social Media Context. *Findings of the Association for Computational Linguistics: EMNLP*, 2020. P.1352–1374. URL: <https://aclanthology.org/2020.findings-emnlp.122/> (дата звернення: 24.03.2025).

33. Закон України «Про основні засади забезпечення кібербезпеки України» № 2145-VIII від 05.10.2017, Верховна Рада України. URL: <https://zakon.rada.gov.ua/laws/show/2145-19#Text> (дата звернення: 24.03.2025).

34. Кодекс України про адміністративні правопорушення, Верховна Рада України. URL: <https://zakon.rada.gov.ua/laws/show/80731-10#n4218> (дата звернення: 26.03.2025).

35. Закон України «Про заборону пропаганди російського нацистського тоталітарного режиму, збройної агресії Російської Федерації як держави-терориста проти України, символіки воєнного вторгнення російського нацистського тоталітарного режиму в Україну», *Ligazakon*. URL: <https://zakon.rada.gov.ua/laws/show/2265-20#n48> (дата звернення: 05.04.2025).

36. Abrar M. F., Khan M. S., Khan I., ElAffendi M., Ahmad S. Towards Fake News Detection: A Multivocal Literature Review of Credibility Factors in Online News Stories and Analysis Using Analytical Hierarchical Process. *Electronics* 2023,

12(15), 3280. URL: <https://doi.org/10.3390/electronics12153280> (дата звернення: 15.04.2025).

37. Молчанова М.О., Бармак О.В. Метод нейромережевого виявлення прийомів пропаганди за маркерами з візуальною інтерпретацією прийнятих рішень. *Науковий журнал «Computer Science and Applied Mathematics»*. 2024. №2. С. 72-82. URL: <https://doi.org/10.26661/2786-6254-2024-2-08> (дата звернення: 16.04.2025).

38. Дурняк Б. В., Музика Д. В, Сабат В. І. Стеганографічні методи захисту документів: монографія. Львів: Українська академія друкарства. 2014. 160 с.

39. Kay See Tan, Yi-Chen Yeh, Prasad S Adusumilli, William D Travis. Quantifying Interrater Agreement and Reliability Between Thoracic Pathologists: Paradoxical Behavior of Cohen's Kappa in the Presence of a High Prevalence of the Histopathologic Feature in Lung Cancer *JTO Clinical and Research Reports*. 2024, Vol. 5, no 1. 100618 URL: <https://doi.org/10.1016/j.jtocrr.2023.100618> (дата звернення: 20.04.2025).

40. Senkivsky V., Sikora L., Lysa N., Kudriashova A., Pikh I. Fuzzy System for the Quality Assessment of Educational Multimedia Edition Design. *Appl. Sci*. 2025. 15(8), 4415. URL: <https://doi.org/10.3390/app15084415>. (дата звернення: 22.04.2025).

41. Senkivsky V., Pikh I., Kudriashova A., Senkivska N., Tupyshak L. Models of Factors of the Design Process of Reference and Encyclopedic Book Editions. *Lecture Notes on Data Engineering and Communications Technologies*, 2022, Vol.77, P. 217–229/

42. Pikh I., Senkivsky V., Kudriashova A., Senkivska N. Prognostic Assessment of COVID-19 Vaccination Levels. *Lecture Notes on Data Engineering and Communications technologies*, 2023, Vol.149, P. 246–265

43. Піх І. В., Сеньківський В.М., Сеньківська Н. Є., Калиній І. В. (2017). Теоретичні основи забезпечення якості видавничо-поліграфічних процесів (Частина 4. Прогнозування та забезпечення якості засобами нечіткої логіки)

Наукові записки. 2017. № 1 (54). С. 22-30.

44. Pikh I., Senkivsky V., Babichev S., Senkivska N., Kudriashova A. Modeling of alternatives and defining the best options for websites design // CEUR Workshop Proceedings. 2021. Vol. 2853: P. 259-270.

45. Senkivsky V., Havrysh B., Holubnyk T., Snihur N. Design and optimization of multi-page publications for high-quality rendering // CEUR Workshop Proceedings. 2023. Vol. 3373: P. 302-319.

46. Дурняк Б. В. Піх І. В. Сеньківський В. М. Теоретичні основи інформаційної концепції формування та оцінювання якості видавничо-поліграфічних процесів: монографія. Львів: Українська академія друкарства. 2022. 356 с.

47. Senkivsky V., Pikh I., Senkivska N., Hileta I., Lytovchenko O., Petyak Y. Forecasting assessment of printing process quality // Advances in Intelligent Systems and Computing (AISC). 2020. Vol. 1246: Lecture notes in computational intelligence and decision making. 2020 International scientific conference "Intellectual systems of decision-making and problems of computational intelligence ISDMCI'2020. P. 467-479.

48. Беспалько О., Ткачук Р.Л., Андріїв Р.Р. (2023). Дослідження методів захисту інформації веб-сайтів на основі моделей розподілення доступу та моніторингу ідентифікаторів користувача. Зб. тез доп. *VI Всеукр. наук.-практ. конф. молодих учених, студентів і курсантів «Інформаційна безпека та інформаційні технології»*. /м. Львів, 30 листопада 2023 р. Львів: ЛДУБЖД. С. 16-20.

49. Durnyak B., Senkivsky V., Pikh I., Kudriashova A. Modelling of Post-printing Processes and Book Interest Level. Lecture Notes on Data Engineering and Communications Technologies, 2024, Vol. 219, P. 3–27.

50. Pikh I., Senkivsky V., Kudriashova A., Tupyachak L., Andriiv R. Formation of an indicator of the information message reliability level by fuzzy logic toolkit. *CEUR Workshop Proceedings*, 2024, Vol. 3899, P. 99-112.

URL: <https://ceur-ws.org/Vol-3899/paper10.pdf> (дата звернення: 03.05.2025).

51. Roman Andriiv. Modern methods of ensuring information protection in cybersecurity systems using artificial intelligence and Blockchain technology: collective monograph / O. Harasymchuk and others. Kharkiv: TECHNOLOGY CENTER PC, 2025. 132 p.

52. Сеньківський В. М., Піх І. В., Калиній І. В., Сеньківський Н. Ю., Драгомиров М. А. Методологічні основи формування якості програмного забезпечення (Частина 1: Базова модель факторів якості). *Поліграфія та видавнича справа*. 2022. № 2(84), С. 9-21

53. Сеньківський В. М., Піх, І. В., Калиній І. В., Литовченко О. В. Моделювання процесів формування та прогнозування якості книжкових видань: монографія. Львів: Українська академія друкарства, 2023. 272 с.

54. Cosgrove, A. L., Beaty, R. E., Diaz, M. T., & Kenett, Y. N. Age differences in semantic network structure: Acquiring knowledge shapes semantic memory. *Psychology and Aging*, 2023. 38(2), P. 87–102. URL: <https://doi.org/10.1037/pag0000721> (дата звернення: 03.05.2025).

55. Senkivskyu, V. M., Pikh, I. V., Kudriashova, A. V., Senkivska, N. E., Kalynii, I. V. Optimization of the model of reader demand factors for the book. *Printing and Publishing*. 2021. No. 2(84), P. 11-20.

56. Eight effective ways to recognize fake information. URL: <https://www.prostir.ua/?news=8-dijevyh-sposobiv-rozpiznaty-fejkovu-informatsiyu>. (дата звернення: 05.05.2025).

57. Піх І. В., Сеньківський В. М., Теслюк В. М., Цмоць І. Г. Моделі факторів інтенсивності вакцинації від COVID-19 з урахуванням предикатів семантичних мереж. *Наукові записки*. 2022. № 1(64). С. 63-75.

58. Certificate of copyright registration for work No. 41832. Ukraine. Simulation modeling in system analysis using binary comparisons method (computer program). The copyrights are owned by I. V. Hileta, V. M. Senkivskyu, O. V. Melnykov. Registered 17.01.2012.

59. Піх І. В., Сеньківський В. М., Андріїв Р. Р. Оптимізована модель чинників достовірності текстових даних. *Науковий вісник Національного лісотехнічного університету України*. 2024. Том 34 № 4, С. 78-85.

URL: <https://doi.org/10.36930/40340410>. (дата звернення: 10.05.2025).

60. Durnyak B., Pikh I., Senkivsky V. Method of determining the weight of predicates of semantic networks in publishing processes. *Przegląd papierniczy (Polish Paper Review)*. Warszawa. 2018. Vol. 1, Pp. 57-60.

61. Піх І. В., Сеньківська Н. Є., Білик О. З., Сеньківський Н. Ю., Андріїв Т. Р., Андріїв Р.Р. Модель факторів якості тестування програмного забезпечення. *Комп'ютерні технології друкарства. Зб. наук. праць*. 2024, № 1(51). С. 90-108.

62. Novorushchenko T., Alekseiko V., Shvaiko V., Ilchyshyna J., Kuzmin A.. Information system for earth's surface temperature forecasting using machine learning technologies. *Computer Systems and Information Technologies*, 2024. (4) P. 51–58.

URL: <https://doi.org/10.31891/csit-2024-4-7> (дата звернення: 10.05.2025).

63. Піх І. В. Кудряшова А. В. Багатофакторний вибір альтернативних варіантів композиційного оформлення видання на основі лінійного згортання критеріїв. *Наукові записки*. 2017. № 2 (55). С. 41-46.

64. Elsisy M. A., El Sayed M. A., Abo-Elnaga Y. A novel algorithm for generating Pareto frontier of bi-level multi-objective rough nonlinear programming problem. *Ain Shams Engineering Journal* 2021. Vol. 12, no 2 P. 2125-2133. URL: <https://doi.org/10.1016/j.asej.2020.11.006> (дата звернення: 12.05.2025).

65. Піх І. В., Дурняк Б. В., Сеньківський В. М., Голубник Т. С. Інформаційні технології формування якості книжкових видань: монографія Львів: НВЛПТ УАД, 2017. 308 с.

66. Mahajan, Sumati, Abhishek Chauhan, and S. K. Gupta. On Pareto optimality using novel goal programming approach for fully intuitionistic fuzzy multiobjective quadratic problems. *Expert Systems with Applications* 2024. Vol. 243. 122816. URL: <https://doi.org/10.1016/j.eswa.2023.122816> (дата звернення: 16.05.2025).

67. Kumar I., Kapil K., Ghosh I. Reliability Estimation in Inverse Pareto

Distribution Using Progressively First Failure Censored Data. *American Journal of Mathematical and Management Sciences* 2023. Vol. 42, no 2, P. 126-147
URL: <https://doi.org/10.1080/01966324.2022.2163441> (дата звернення: 16.05.2025).

68. Піх І.В. Сеньківський В.М., Андріїв Р.Р. Проектування та розрахунок альтернативних варіантів реалізації технологічних процесів. *Технологія і техніка друкарства*. 2015 № 2 (48). С. 55-62. URL: [https://doi.org/10.20535/2077-7264.2\(48\).2015.48030](https://doi.org/10.20535/2077-7264.2(48).2015.48030). (дата звернення: 17.05.2025).

69. Піх І. В., Андріїв Р. Р. Вибір альтернативної системи для формування навчальних електронних ресурсів. *Поліграфія і видавнича справа*. 2015. № 1 (69), С. 69-76.

70. Сеньківський В. М., Кудряшова А. В., Козак Р. О. Інформаційна технологія формування якості редакційно-видавничого процесу: монографія. Львів: Українська академія друкарства, 2019. 272 с.

71. Pikh I., Senkivskyu V., Sikora L., Lysa N., Kudriashova A. Automated system for evaluating alternatives for developing innovative IT projects. *Appl. Sci.* 2025, 15(3), 1167. URL: <https://doi.org/10.3390/app15031167> (дата звернення: 20.05.2025).

72. Joshi N., Sudeep R., Bapat R., Sengupta R. N. Optimal estimation of reliability parameter for inverse Pareto distribution with application to insurance data. *International Journal of Quality & Reliability Management*, 2024. Vol. 41, no7. P.1811–1837. URL: <https://doi.org/10.1108/IJQRM-06-2023-0213> (дата звернення: 20.05.2025).

73. Піх І.В., Білик О.З., Сеньківський Н.Ю., Андріїв Р.Р., Браташ С.П. (2024). Багатофакторний вибір альтернативних варіантів процесу тестування ПЗ. *Вісник Вінницького політехнічного інституту*. № 3. С. 78-85.

URL: <https://doi.org/10.31649/1997-9266-2024-174-3-78-85> (дата звернення: 23.05.2025).

74. Andriiv, R., Senkivskyu, V. Information technology for predicting the reliability level of text messages. *Computer Systems and Information Technologies*.

2025. (2), 65–71. URL: <https://doi.org/10.31891/csit-2025-2-7>. (дата звернення: 23.05.2025).

75. Dumitrescu C., Ciotirnae P., Vizitiu C. Fuzzy logic for intelligent control system using soft computing applications. *Sensors*, 2021. 21(8). 2617.

URL: <https://doi.org/10.3390/s21082617> (дата звернення: 25.05.2025).

76. Jing J., Wu H., Sun J., Fang X., Zhang H. Multimodal fake news detection via progressive fusion networks. *Information Processing & Management*, 2023. 60(1) 103120. URL: <https://doi.org/10.1016/j.ipm.2022.103120>. (дата звернення: 24.05.2025).

77. Желдак Т. А.; Коряшкіна Л. С.; Ус С. А. Нечіткі множини в системах управління та прийняття рішень: навчальний посібник. Національний технічний університет "Дніпровська політехніка". Київ : НТУ ДП, 2020. 386 с.

78. Chonghui Z., Dong X., Shouzhen Z., Carlos L.-A. Dual consistency-driven group decision making method based on fuzzy preference relation *Expert Systems with Applications*, 2024. Vol. 238. 122228.

URL: <https://doi.org/10.1016/j.eswa.2023.122228> (дата звернення: 24.05.2025).

79. Піх, І., Сеньківський, В., Сеньківська, Н., Калиній, І., Білик, О. Прогнозування якості програмного забезпечення на основі нечіткої логіки. (Частина 1. Функції належності лінгвістичних змінних). *Вимірjuвальна та обчислювальна техніка в технологічних процесах. Хмельницький національний університет (Україна)*. 2024. № 4. С. 7–17. URL: <https://doi.org/10.31891/2219-9365-2024-80-1> (дата звернення: 25.05.2025).

80. Кудряшова А. В., Сеньківський В. М. Опрацювання функцій належності лінгвістичних змінних проєктування післядрукарських процесів (частина 2. Візуалізація значень функцій належності). *Поліграфія і видавнича справа*. 2020. № 1 (79), С. 30-41.

81. Senkivskyuy V., Pikh I., Babichev S., Havenko S. Model of Logical Inference and Membership Functions of Factors for the Printing Process Quality Formation // *Advances in Intelligent Systems and Computing (AISC)*. 2020. Vol. 1020: Lecture

notes in computational intelligence and decision making. Proceedings of the XV International scientific conference Intellectual systems of decision making and problems of computational intelligence (ISDMCI'2019), (Ukraine, May 21-25, 2019). P. 609-621.

82. Kudriashova A., Pikh I., Senkivskyi V., Merenych Y. Evaluation of prototyping methods for interactive virtual systems based on fuzzy preference relation. *Eastern-European Journal of Enterprise Technologies*. 2024. No. 5/4 (131). P. 71-81. URL: <https://doi.org/10.15587/1729-4061.2024.313099>. (дата звернення: 25.05.2025).

83. Andriiv, R., Pikh, I. Method for ranking the reliability factors of text messages. *Computer Systems and Information Technologies*. 2025. (1), P.178–184. URL: <https://doi.org/10.31891/csit-2025-1-21> (дата звернення: 26.05.2025).

84. Дурняк Б.В., Сеньківський В. М., Піх І. В. Формування якості процесу проектування видань на основі матриць знань і нечітких логічних рівнянь *Поліграфія і видавнича справа*. 2017. № 1 (73). С. 11-18.

85. Шаховська Н. Б., Литвин В. В. Проектування інформаційних систем: навчальний посібник. Львів: «Магнолія-2006», 2011. 380 с.

86. Сеньківський В. М., Кудряшова А. В. Моделі інформаційної технології проектування післядрукарських процесів: монографія. Львів: УАД, 2022. 204 с.

87. Сеньківський В. М., Піх І. В., Кудряшова А. В. Теоретичні основи інформаційної технології прогностичного оцінювання якості проектування післядрукарських процесів. *New information technologies, simulation and automation: Monograph* / Velychko V., Voinova S., Granyak V., et al; Editor-in-Chief Kotlyk S. Iowa State University Digital Press, 2022. 729 p. (P. 44–138).

88. Кудряшова А. В. Багатофакторний вибір альтернативних варіантів проектування післядрукарських процесів на основі лінійного згортання критеріїв. *Поліграфія і видавнича справа*. 2019. № 2 (78). С. 45–50.

80. Кудряшова А. В., Литовченко Н. М. Формування інтегрального показника якості процесу структурування видання. *Поліграфія і видавнича справа*.

Львів, 2018. № 1 (75). С. 82–89.

90. Осінчук О. І., Литовченко О. В., Калиній І. В. Формування інтегральних показників якості планування та художньо-технічного оформлення видань засобами теорії нечітких множин. *Моделювання та інформаційні технології*. 2018. Вип. 85. С. 137–142.

91. Sahoo S. R., Gupta B. B. Multiple features based approach for automatic fake news detection on social networks using deep learning. *Applied Soft Computing*, 2021. Vol. 100. 106983. URL: <https://doi.org/10.1016/j.asoc.2020.106983> (дата звернення: 26.05.2025).

92. Aïmeur E., Amri S., Brassard G. Fake news, disinformation and misinformation in social media: a review. *Social Network Analysis and Mining* 13, 2023, Vol. 13. No 30. URL: <https://doi.org/10.1007/s13278-023-01028-5>. (дата звернення: 27.05.2025).

93. Сеньківський В. М., Кудряшова А. В. Моделі інформаційної технології проєктування післядрукарських процесів: монографія. Львів: УАД, 2022. 204 с.

94. Walters W. The Effectiveness of Software Designed to Detect AI-Generated Writing: A Comparison of 16 AI Text Detectors. *Open Information Science*, 2023. Vol. 7, No. 1, 2023, 20220158. URL: <https://doi.org/10.1515/opis-2022-0158> (дата звернення: 29.05.2025).

95. Duggineni, S. Impact of controls on data integrity and information systems. *Science and Technology*, 2023. Vol. 13. No. 2. P.29-35. URL: <https://doi.org/10.5923/j.scit.20231302.04> (дата звернення: 29.05.2025).

96. Яшина О. С., Пісклова Т. С. Проєктування інформаційних систем: навч. посібник. Харків: Національний аерокосмічний університет ім. М. Є. Жуковського «Харків. авіац. інститут», 2024. 68 с.

ДОДАТОК А. ПЕРЕЛІК ПУБЛІКАЦІЙ ЗДОБУВАЧА

Статті у наукових виданнях, включених до Переліку наукових фахових видань України:

1. Піх І.В. Сеньківський В.М., Андріїв Р.Р. Проектування та розрахунок альтернативних варіантів реалізації технологічних процесів. *Технологія і техніка друкарства*. 2015 № 2 (48). С. 55-62. URL: [https://doi.org/10.20535/2077-7264.2\(48\).2015.48030](https://doi.org/10.20535/2077-7264.2(48).2015.48030).

2. Піх І. В., Сеньківський В. М., Андріїв Р. Р. Оптимізована модель чинників достовірності текстових даних. *Науковий вісник Національного лісотехнічного університету України*. 2024. Том 34 № 4. С. 78-85. URL: <https://doi.org/10.36930/40340410>. (Index Copernicus)

3. Піх І.В., Білик О.З., Сеньківський Н. Ю., Андріїв Р.Р. Браташ С.П. Багатофакторний вибір альтернативних варіантів процесу тестування ПЗ. *Вісник Вінницького політехнічного інституту*. 2024. № 3. С. 78-85. URL: <https://doi.org/10.31649/1997-9266-2024-174-3-78-85>. (Index Copernicus)

4. Гарасимчук О., Наконечний Ю., Луковський Т., Р. Андріїв Р., Наконечний Т. Захищене зберігання даних із використанням технології блокчейн ETHEREUM. *Ukrainian Information Security Research Journal*. 2024. Том 26 № 1. С. 213-220. URL: <https://doi.org/10.18372/2410-7840.26.18845>.

5. Andriiv, R., Pikh, I. Method for ranking the reliability factors of text messages. *Computer Systems and Information Technologies*. 2025 (1), 178–184. URL: <https://doi.org/10.31891/csit-2025-1-21>.

6. Andriiv, R., Senkivsky, V. Information technology for predicting the reliability level of text messages. *Computer Systems and Information Technologies*. 2025. (2), 65–71. URL: <https://doi.org/10.31891/csit-2025-2-7>.

Монографії (розділи у колективних монографіях):

7. Andriiv Roman. Modern methods of ensuring information protection in cybersecurity systems using artificial intelligence and Blockchain technology:

collective monograph / O. Harasymchuk and others. Kharkiv: TECHNOLOGY CENTER PC, 2025. 132 p. (ISBN-13 (15) 978-617-8360-12-2) URL: <http://doi.org/10.15587/978-617-8360-12-2>

Публікації, які засвідчують апробацію матеріалів дисертації:

8. Піх І., Андріїв Р. Ідентифікація факторів впливу на якість проектування електронних видань. *Науково-технічна конференція професорсько-викладацького складу, наукових працівників і аспірантів: тези доповідей*, м. Львів, 16 лютого – 20 лютого 2015 р. / УАД Львів, 2015. С.118.

9. Піх І., Андріїв Р. Об'єктно-орієнтована модель процесу проектування електронних видань. *Науково-технічна конференція професорсько-викладацького складу, наукових працівників і аспірантів: тези доповідей*, м.Львів, 16 лютого – 19 лютого 2016 р. / УАД, Львів, 2016. С. 145.

10. Беспалько О., Ткачук Р., Андріїв Р. Дослідження методів захисту інформації веб-сайтів на основі моделей розподілення доступу та моніторингу ідентифікаторів користувача. Зб. тез доп. *VI Всеукр. наук.-практ. конф. молодих учених, студентів і курсантів «Інформаційна безпека та інформаційні технології»*, м.Львів, 30 листопада 2023 р. / ЛДУБЖД Львів, 2023. С. 16-20.

11. Pikh I., Senkivskyi V., Kudriashova A., Turychak L., Andriiv R. Formation of an indicator of the information message reliability level by fuzzy logic toolkit. *CEUR Workshop Proceedings*, 2024, Vol. 3899, pp. 99-112. URL: <https://ceur-ws.org/Vol-3899/paper10.pdf> (індексована в наукометричній базі Scopus)

12. Senkivskyi V., Pikh I., Kudriashova A., Andriiv R., Senkivskyi N.. Evaluation of methods for intellectual analysis of manipulative content in the information space. *CEUR Workshop Proceedings*, 2025, Vol. 3963, pp. 139-155. URL: <https://ceur-ws.org/Vol-3963/paper12.pdf> (індексована в наукометричній базі Scopus)

Публікації, які додатково відображають наукові результати дисертації:

13. Сеньківський В. М., Піх І. В., Андріїв Р. Р. Синтез моделі факторів композиційного оформлення іміджевої презентації. *Поліграфія і видавнича справа*. 2011. № 4 (56), С. 117-124. (Index Copernicus)

14. Андріїв Р. Р., Піх І. В. Вибір альтернативної системи для формування навчальних електронних ресурсів. *Поліграфія і видавнича справа*. 2015. № 1 (69), С. 69-76. (Index Copernicus)

15. Піх І. В., Сеньківська Н. Є., Білик О. З., Сеньківський Н. Ю., Андріїв Т. Р., Андріїв Р. Р. Модель факторів якості тестування програмного забезпечення. *Комп'ютерні технології друкарства*. 2024 № 1(51). С. 90-108. URL: <http://doi.org/10.32403/2411-9210-2024-1-51-90-108>

16. Піх І. В., Андріїв Р. Р., Коцік О. Ю. (2024). Інноваційні підходи до розробки та впровадження стартапів. *Науково-технічна конференція професорсько-викладацького складу, наукових працівників і аспірантів: тези доповідей*. м. Львів, 5 лютого – 9 лютого 2024 р./ УАД Львів, 2024. С.186

ДОДАТОК Б. ДОВІДКИ ПРО ВИКОРИСТАННЯ В НАВЧАЛЬНОМУ ПРОЦЕСІ РЕЗУЛЬТАТІВ ДИСЕРТАЦІЙНОЇ РОБОТИ

Директор Інституту поліграфії
та медійних технологій

НУ «Львівська політехніка»

 Ярослав УГРИН
«20» 2025 р.

ДОВІДКА

**Про використання матеріалів дисертації Андріїва Р.Р.
«Інформаційна технологія оцінювання достовірності
текстових повідомлень»**

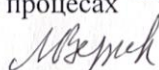
У сучасному цифровому середовищі проблема перевірки достовірності текстових інформаційних повідомлень набуває особливої гостроти. Швидке зростання обсягів онлайн-контенту, поширення соціальних мереж та легкість доступу до каналів масової комунікації сприяють не лише оперативному обміну знаннями, а й масштабному розповсюдженню недостовірної інформації. Для ефективного вирішення цих викликів необхідно застосовувати міждисциплінарний підхід, що поєднує інструменти обробки природної мови та системного моделювання й уможливорює об'єктивну автоматизовану оцінку достовірності даних в сучасних умовах.

Вказане знайшло відображення у навчальному процесі Інституту поліграфії та медійних технологій під час викладання дисциплін: «Моделювання інформаційних систем і процесів», «Моделі і методи прийняття рішень», «Проектування інформаційних систем», «Методи проектування видавничих мультимедійних систем».

Під час вивчення наведених дисциплін використовуються алгоритми, моделі та методи, розроблені і реалізовані яких виконано у дисертаційній роботі Андріїва Р.Р. «Інформаційна технологія оцінювання достовірності текстових повідомлень». До них належать графічно орієнтовані моделі відображення відношень між факторами достовірності повідомлень, оптимізовані моделі чинників впливу на достовірність даних, моделі логічного виведення, нечіткі бази знань та системи нечітких логічних рівнянь, автоматизована система визначення вірогідності контенту отриманих новин, що у підсумку формують інформаційну технологію оцінювання достовірності текстових повідомлень.

Матеріали дисертації використовуються під час визначення тем та змісту бакалаврських і магістерських кваліфікаційних робіт студентів інженерно-технологічних спеціальностей.

Завідувач кафедри комп'ютерних технологій
у видавничо-поліграфічних процесах
д. т. н., професор



Михайло ВЕРХОЛА





ТВЕРДЖУЮ

Проректор університету із навчально-методичної роботи
Львівського державного університету
безпеки життєдіяльності
полковник служби цивільного захисту
Олександр ПРИДАТКО
10.01.2025 2025 р.

АКТ ВПРОВАДЖЕННЯ

результатів дисертаційного дослідження АНДРІЙВА Романа Ростиславовича
за темою «Інформаційна технологія оцінювання достовірності текстових повідомлень»
в освітній процес кафедри управління інформаційною безпекою

Комісія у складі – голови начальника кафедри управління інформаційною безпекою, д.т.н., професора, полковника служби цивільного захисту ТКАЧУКА Ростислава Львовича та членів: доцента кафедри управління інформаційною безпекою, к.т.н., доцента, підполковника служби цивільного захисту ІВАНУСИ Андрія Івановича і доцента кафедри управління інформаційною безпекою, к.т.н., доцента ПОЛОТАЯ Ореста Івановича.

Цим актом засвідчується, що результати дисертаційного дослідження АНДРІЙВА Романа Ростиславовича на тему «Інформаційна технологія оцінювання достовірності текстових повідомлень» впроваджено в освітній процес Львівського державного університету безпеки життєдіяльності. Елементи наукового дослідження інтегровані у зміст навчальних дисциплін «Політика інформаційної безпеки в організаціях» та «Менеджмент інформаційної безпеки в системі цивільного захисту», «Системи охорони державної таємниці», що викладаються здобувачам першого (бакалаврського) рівня вищої освіти, а також у дисципліну «Інтегровані системи санкціонованого доступу до інформації», що реалізується в межах підготовки здобувачів другого (магістерського) рівня за спеціальністю F5 «Кибербезпека та захист інформації».

У межах вказаних освітніх компонент використовуються основні положення дисертації, зокрема ієрархічно структуровані, упорядковані за ваговими коефіцієнтами та оптимізовані графічні моделі вагових значень пріоритетного впливу факторів на достовірність інформаційних повідомлень, пов'язаних у свою чергу з безпекою навколишнього середовища; модель логічного виведення у вигляді багаторівневої ієрархічної структури, що відтворює послідовність і логіку процесу формування інтегрального показника рівня достовірності повідомлень, інформаційна технологія прогнозного оцінювання вірогідності повідомлень на основі методів теорії нечітких множин засобів нечіткого моделювання.

Зазначені результати дисертаційного дослідження впроваджено у формі навчально-методичних матеріалів, оновлених робочих навчальних програм (силабусів) освітніх компонент, практичних кейсів, орієнтованих на набуття професійних компетентностей у сфері оцінювання достовірності текстових повідомлень.

Голова комісії:

начальник кафедри управління
інформаційною безпекою, д.т.н., професор

Ростислав ТКАЧУК

Члени комісії:

доцент кафедри управління
інформаційною безпекою, к.т.н., доцент

Андрій ІВАНУСА

доцент кафедри управління
інформаційною безпекою, к.т.н., доцент

Орест ПОЛОТАЙ

ДОДАТОК В. АКТ ВИПРОБУВАНЬ ПРОГРАМНОГО ЗАСТОСУНКУ ВИЗНАЧЕННЯ ПРОГНОЗНОГО РІВНЯ ДТІП

АКТ впровадження результатів дисертаційної роботи Андрієва Романа Ростиславовича

Видавничо-поліграфічні процеси в умовах сучасного динамічного розвитку виробничих відносин, комп'ютеризації виробництва, активного впровадження інформаційних технологій та обмеження енергетичних ресурсів потребують суттєвої переорієнтації й перегляду пріоритетів щодо забезпечення якості, достовірності та змістової цілісності текстових інформаційних повідомлень. Вирішення зазначених проблем вимагає впровадження комплексного підходу, що поєднує методи аналізу текстових даних, засоби штучного інтелекту, алгоритми обчислювальної лінгвістики та інструменти автоматизованого контролю якості, забезпечуючи високий рівень надійності інформації в умовах сучасного інформаційного середовища.

Одними із вагомих та актуальних результатів дисертаційної роботи Андрієва Р.Р. «Інформаційна технологія оцінювання достовірності текстових повідомлень» є розроблені моделі та засоби організаційного, технічного та соціального спрямування, що обумовили створення інформаційної технології та її складової – автоматизованої системи, що в комплексі забезпечують ефективне оцінювання вірогідності сформованих у видавничих структурах та отриманих користувачами текстових інформаційних повідомлень.

Практичне значення для видавничої організації має розроблена здобувачем автоматизована система оцінювання рівня достовірності текстових повідомлень, яка дозволяє здійснювати перевірку редакційного контенту на вірогідність із застосуванням спеціалізованого програмного забезпечення. Використання цієї системи сприяє вибору оптимальних варіантів подання текстової інформації, забезпечує підвищення якості редакційного аналізу, а також мінімізацію ризиків поширення недостовірних або упереджених даних, що, у свою чергу, підвищує рівень довіри до змісту інформаційних повідомлень.

Директор друкарні ТОВ
«Видавничий Дім «Високий замок»»



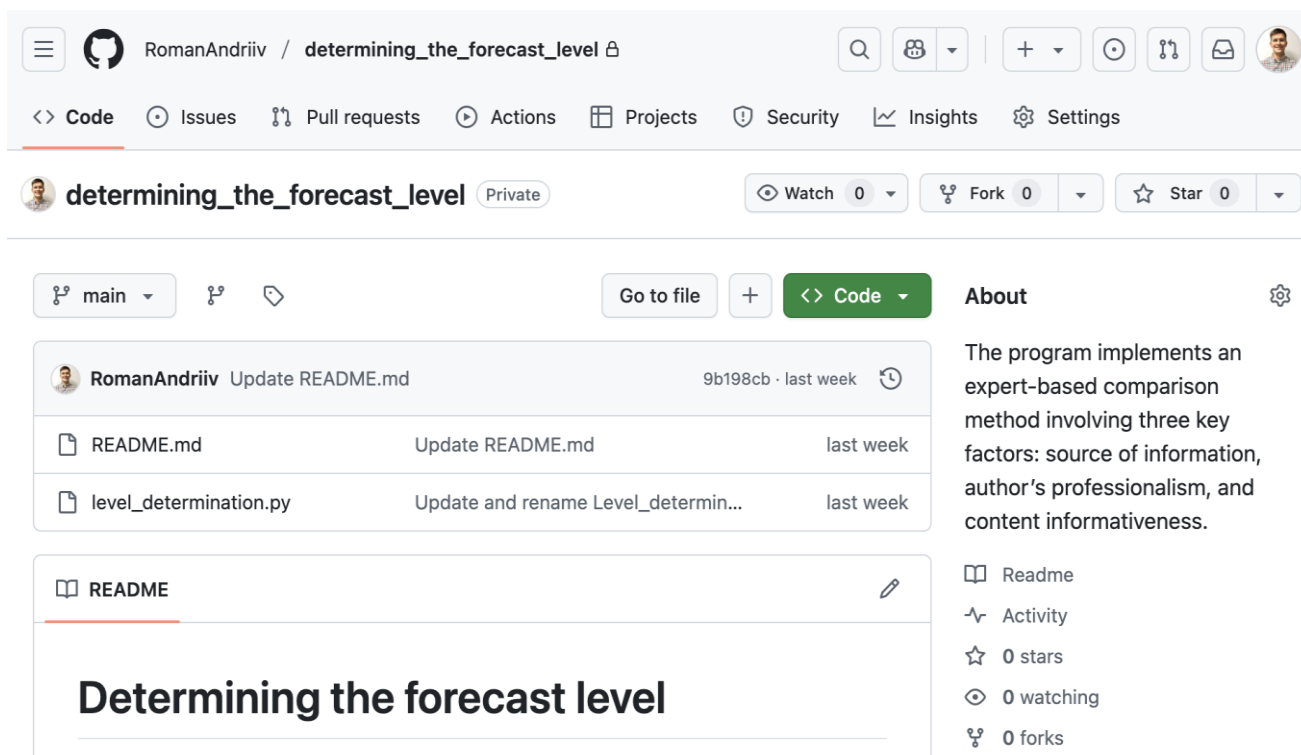
Салагай П.Н.

05.02.2025 р.

ДОДАТОК Г. ПРОГРАМНИЙ КОД

Вихідний код, використаний у дослідженні, доступний у репозиторії

GitHub: https://github.com/RomanAndriiv/determining_the_forecast_level (дата звернення: 14.07.2025). На рисунку нижче зображено скріншот репозиторію.



У файлі «level_determination.py» міститься програмний код для запуску графічного інтерфейсу користувача, що призначений для визначення прогнозного рівня достовірності текстових повідомлень.

Програма реалізує метод експертного порівняння трьох факторів: джерело інформації, професіоналізм автора, інформативність контенту. На основі введених даних відбувається побудова відповідних матриць, обчислення функцій належності та визначення рівня достовірності повідомлення як одного з трьох класів: «низький», «середній» або «високий».

Крім основної функціональності, у програмі реалізовані можливості очищення вхідних даних та завершення роботи програми. Інтерфейс створено з використанням бібліотеки Tkinter.