



В. В. Пасічник, М. В. Яромич

Національний університет "Львівська політехніка", м. Львів, Україна

ІНФОРМАЦІЙНА ТЕХНОЛОГІЯ АВТОМАТИЗОВАНОГО РЕФЕРУВАННЯ НАУКОВИХ ТЕКСТІВ ЗАСОБАМИ ВЕЛИКИХ МОВНИХ МОДЕЛЕЙ

Наведено комплексну інформаційну технологію автоматизованого реферування наукових текстів проблемної галузі "штучний інтелект", що інтегрує екстрактивні методи виділення змістових фрагментів, абстрактивні моделі узагальнення на основі сучасних трансформерних архітектур та механізми автоматизованого оцінювання якості отриманих результатів. Технологія реалізована як цілісний багаторівневий цикл інтелектуального оброблення тексту, який охоплює попередню структуру його сегментацію й очищення документа, формування компактних змістових подань, багатоступеневе генеративне узагальнення та верифікацію підсумків за формальними і семантичними метриками. Проведено три експериментальні спостереження із застосуванням моделей GPT-5, Grok-4 та Gemini 2.5 на англійськомовному, німецькомовному та українськомовному наукових текстах. Отримані результати продемонстрували високу якість збереження логічної структури наукових праць, міждисциплінарних концептуальних зв'язків і термінологічної коректності, а також підтвердили багатомовність і доменну адаптивність розробленої інформаційної технології. Модель GPT-5 забезпечила найбільшу глибину концептуального узагальнення та інтеграцію наукових ідей у цілісні когнітивні структури, водночас як модель Grok-4 продемонструвала чутливість до філософсько-аналітичного стилю німецьких наукових текстів і складних теоретичних аргументацій, а модель Gemini 2.5 виявила високу точність і стилістичну узгодженість під час роботи з українською технічною мовою та спеціалізованою науковою термінологією. Запропонована інформаційна технологія забезпечує формування зв'язних, інформативних і фактологічно коректних рефератів для різних мов і наукових дисциплін, підтримує адаптацію до різних предметних галузей та створює підґрунтя для її практичного застосування в системах інтелектуального аналізу знань, наукових репозиторіях, освітніх платформах, цифрових бібліотеках та сервісах автоматизованого огляду наукової літератури, спрямованих на підтримку дослідницької діяльності, міждисциплінарної інтеграції та пришвидшення наукової комунікації.

Ключові слова: автоматизоване реферування; великі мовні моделі; екстрактивне скорочення тексту; абстрактивне узагальнення тексту; багатомовні наукові тексти.

Вступ / Introduction

Для сучасного наукового простору характерне стрімке зростання кількості публікацій, особливо у сфері штучного інтелекту, де темпи появи нових досліджень випереджають можливості їх ручного опрацювання. Збільшення обсягів інформації призводить до того, що ефективне засвоєння наукових матеріалів стає складнішим, а процес пошуку релевантної літератури – значно тривалішим. Дослідникам важко підтримувати актуальність знань, відстежувати нові тренди та формувати повноцінні огляди літератури, тому зростає потреба в інструментах, здатних автоматично створювати якісні підсумки наукових текстів. Проблема ускладнюється ще й тим, що академічні публікації у сфері штучного інтелекту містять складні аргументативні структури, формальні визначення, експериментальні результати та міждисциплінарні концепти, що робить їх значно важчими для автоматизованого аналізу порівняно з текстами загального призначення [24]. Класичні методи скорочення текстів, засновані на статистичних підходах

або поверхневому аналізі структури документа, демонструють обмежену придатність у таких умовах, оскільки не здатні повноцінно відтворити наукову логіку та семантичні зв'язки [5].

Поява великих мовних моделей істотно розширила потенціал автоматизованого реферування [23]. Моделі нового покоління здатні працювати з великими обсягами тексту, утримувати довготермінові залежності, виконувати абстрактивне узагальнення та генерувати реферати, які за стилем і структурою наближені до людських. Проте, дослідження показують, що навіть найсучасніші моделі схильні до фактологічних помилок, опущення важливих фрагментів або навпаки – до надмірного узагальнення, що змінює смислову точність первинного документа [17]. Це вимагає створення спеціалізованих технологій, які не тільки генерують підсумки, але й враховують специфіку академічних текстів, а також забезпечують механізми внутрішньої перевірки якості результату.

Значним викликом є також проблема багатомовності. Більшість досліджень у сфері summarization орієнто-

© Copyright 2026 by the author(s)

Пасічник Володимир Володимирович, д-р техн. наук, професор, кафедра інформаційних систем та мереж.

Email: vpasichnyk@gmail.com; <https://orcid.org/0000-0002-5231-6395>

Яромич Максим Віталійович, аспірант, кафедра прикладної лінгвістики. Email: yamax0312@gmail.com;

<https://orcid.org/0009-0005-3299-6695>

Цитування за ДСТУ: Пасічник В. В., Яромич М. В. Інформаційна технологія автоматизованого реферування наукових текстів засобами великих мовних моделей. Науковий вісник НЛТУ України. 2026, т. 36, № 1. С. 105–117.

Citation APA: Pasichnyk, V. V., & Yaromych, M. V. (2026). Information technology of automated reference of scientific texts using large language models. *Scientific Bulletin of UNFU*, 36(1), 105–117. <https://doi.org/10.36930/40360112>

вані переважно на англomовні корпуси, тоді як світова наука функціонує десятками мов, і потреба у підтримці багатомовних корпусів тільки зростає. Сучасні трансформерні моделі демонструють потенціал у міжмовному узагальненні. Однак, якість їхніх результатів значно залежить від мови, структури тексту та специфіки предметної галузі [14]. Це зумовлює потребу в технологіях, здатних стабільно працювати з науковими текстами різних мов, зберігаючи точність і повноту змісту.

Окремим аспектом актуальності тематики дослідження є нагальна потреба поєднання екстрактивних та абстрактивних методів. Екстрактивні підходи забезпечують збереження фактологічної точності, але часто втрачають когерентність, тоді як абстрактивні моделі здатні створювати більш природні та зв'язні тексти. Проте, іноді жертвують точністю термінології. Дослідження свідчать, що комбіновані підходи здатні істотно покращити результати, оскільки дають можливість узгодити достовірність з когерентністю. Використання графових моделей, таких як TextRank та його модифікації, також демонструє ефективність у збереженні ключових семантичних вузлів наукового тексту [1], що підсилює обґрунтованість інтеграції цих підходів у єдину технологічну систему.

У межах наукового аналізу дедалі більшої уваги набуває питання достовірності, пояснюваності та відтворюваності результатів автоматизованого реферування. Академічні тексти вимагають точності та прозорості, тому результати, згенеровані моделлю, мають підлягати автоматизованому оцінюванню за допомогою формальних метрик, таких як ROUGE, BLEU чи BERTScore. У недавніх роботах наголошено, що навіть провідні моделі демонструють систематичні перекоси у відтворенні наукової структури, що вимагає ретельних механізмів валідації [32]. Зокрема, встановлено, що LLM можуть несвідомо посилювати другорядні частини тексту або трансформувати логічні зв'язки, що критично впливає на наукові висновки та об'єктивність підсумків [26].

Усе це вказує на важливість створення комплексної інформаційної технології, яка реалізує повний цикл автоматизованого реферування: від попереднього оброблення тексту до редакційно-оцінювального контролю результатів. Така система має бути здатною працювати з науковими текстами різних предметних галузей, підтримувати багатомовність, комбінувати екстрактивні та абстрактивні алгоритми, забезпечувати стабільну якість підсумків та давати змогу інтегрувати різні моделі. Стрімке зростання обсягів наукової літератури, зокрема – в галузі штучного інтелекту, робить розроблення подібної технології не тільки актуальною, але й необхідною для ефективного функціонування сучасної науки. Автоматизоване реферування текстової інформації здатне зменшити аналітичне навантаження на дослідників, пришвидшити процедуру формування аналізу літератури, сприяти інтеграції між мовами та дисциплінами й стати основою для нових когнітивних систем підтримки наукових досліджень.

Об'єкт дослідження – автоматизоване реферування наукових текстів із використанням прикладної лінгвістики та великих мовних моделей.

Предмет дослідження – інформаційна технологія автоматизованого реферування наукових текстів, що поєднує екстрактивні та абстрактивні методи оброблення тексту, багатомовне генерування рефератів і механізми

автоматизованого оцінювання якості рефератів на основі формальних та семантичних метрик.

Мета роботи – розробити та експериментально перевірити комплексну інформаційну технологію автоматизованого реферування наукових текстів проблемної галузі "штучний інтелект" із використанням великих мовних моделей, здатну забезпечити змістову точність, логічну узгодженість і багатомовність рефератів, а також обґрунтувати можливості її застосування в наукових, освітніх та аналітичних інформаційних системах.

Для досягнення зазначеної мети визначено такі *основні завдання дослідження*:

- обґрунтувати вибір технологічної схеми автоматизованого реферування, яка поєднувала б екстрактивні та абстрактивні підходи, що дасть змогу забезпечити баланс між фактологічною точністю та когерентністю згенерованих рефератів;
- визначити та формалізувати основні етапи інформаційної технології автоматизованого реферування, що дасть можливість систематизувати процес опрацювання наукових текстів від попередньої сегментації до фінального узагальнення;
- розробити підхід до використання великих мовних моделей у процесі абстрактивного узагальнення наукових текстів, що забезпечить збереження логічної структури, термінологічної коректності та академічного стилю;
- реалізувати механізми автоматизованого оцінювання якості результатів реферування з використанням формальних і семантичних метрик, які допоможуть об'єктивно порівнювати згенеровані реферати з авторськими анотаціями;
- провести експериментальну перевірку запропонованої технології з використанням великих мовних моделей GPT-5, Grok-4 та Gemini 2.5, які уможливають аналіз ефективності технології в умовах багатомовних наукових корпусів;
- дослідити вплив мовних і стилістичних особливостей на якість автоматизованого реферування, які спонукають до виявлення переваг і обмежень окремих моделей у різних мовних традиціях;
- проаналізувати результати експериментальних спостережень та узагальнити закономірності функціонування розробленої інформаційної технології, які приведуть до визначення перспектив її подальшого використання та розвитку у сфері інтелектуального аналізу наукових текстів.

Аналіз останніх досліджень та публікацій. Дослідження автоматизованого реферування наукових текстів має тривалу історію та охоплює широкий спектр методів, що еволюціонували від статистичних моделей до сучасних трансформерних архітектур. Ранні роботи у сфері текстового скорочення базувалися переважно на частотному аналізі слів, векторних моделях та метриках важливості термінів. Методи TF-IDF, латентно-семантичний аналіз та прості кластеризатори відіграли важливу роль у формуванні перших інструментів, здатних автоматично виявляти ключові сегменти текстів [11]. Однак, такі підходи мали значні обмеження: вони погано працювали з великими науковими текстами, не враховували дискурсивну структуру документа та втрачали семантичні зв'язки між ідеями [19].

Подальший розвиток пов'язаний зі становленням графових методів. Алгоритми TextRank та LexRank, засновані на обчисленні центральності речень, започаткували нову хвилю у напрямі екстрактивного реферування, демонструючи кращу здатність працювати із структурованими текстами та зберігати ключові тези доку-

мента. Проте, графові методи також виявилися недостатньо ефективними щодо складних наукових текстів, оскільки вони розглядали речення як незалежні вузли, не моделюючи глибини аргументації та логічних переходів [31].

Значного прогресу в автоматизованому створенні підсумків було досягнуто з появою нейронних моделей для узагальнення тексту. Sequence-to-sequence архітектури відкрили шлях до абстрактивного підходу, у якому модель формує підсумок у новій мовній формі, а не просто вибирає речення з оригінального тексту. Проте, класичні seq2seq моделі на базі LSTM та GRU виявилися недостатніми для роботи з довгими документами, оскільки не могли ефективно утримувати контекст великого обсягу [29].

Поява трансформерів стала ключовою точкою у розвитку сучасного автоматизованого реферування. Архітектура Attention дала змогу моделювати далекі зв'язки між частинами тексту, а моделі типу BART, PEGASUS та T5 помітно покращили якість абстрактивного узагальнення. Дослідження показали, що трансформери здатні не тільки передавати основну думку документа, але й зберігати логіку викладу, що особливо важливо для наукових текстів [27]. Однак, навіть ці моделі демонстрували складнощі у роботі зі статтями значного обсягу, через обмеження довжини контексту та ризик семантичного "згортання" інформації [20].

У межах досліджень останніх років значну увагу приділяють питанням достовірності та фактологічної точності генерованих рефератів. Проблема галюцинацій є актуальною особливо у сфері наукових текстів, де спотворення логіки або результатів дослідження може істотно вплинути на розуміння наукових висновків [22]. У низці робіт запропоновано поєднання екстрактивних та абстрактивних методів для зниження ризику втрат важливої інформації та підтримки послідовності викладу [10].

Дослідження мультимовного реферування також набуває дедалі більшої актуальності. Наголошено на тому, що якість підсумків значно залежить від мови та стилістичних особливостей тексту. Багатомовні моделі демонструють кращу узагальнюваність, Проте, нерівномірно працюють із науковими корпусами різних мов, що спонукає до розроблення спеціалізованих стратегій для багатомовних наукових матеріалів [3].

Окремий напрям досліджень стосується опрацювання великих документів. Новітні архітектури, зокрема – Longformer, BigBird та Llama-3-Long, намагаються подолати проблему обмеженої довжини контексту завдяки механізмам розрідженої або ієрархічної уваги. У сучасних роботах доведено, що моделі з розширеним контекстним вікном значно підвищують якість роботи з науковими текстами, оскільки можуть охоплювати повну структуру статті, враховуючи методологію, експерименти та висновки [2].

Дослідження останніх двох років активно спрямовані на оцінювання результатів автоматичного реферування. Стандартні метрики, такі як ROUGE, не завжди адекватно відображають якість наукових підсумків, оскільки орієнтовані на поверхневі збіги. Натомість запропоновано метрики на основі семантичної відповідності, такі як BERTScore та QuestEval, які краще відображають змістову точність і узгодженість із джерелом [28]. Також активно досліджують методи людського оці-

нювання, що демонструють значно більшу надійність у випадку академічних текстів.

Новий напрям, який активно розвивається, – використання LLM (англ. *Large Language Model*) як мета-оцінювачів підсумків. Моделі GPT-4, Claude 3 та Gemini продемонстрували конкурентоспроможні результати у завданнях оцінювання, іноді перевершуючи класичні метрики [21]. Це відкриває шлях до інтеграції внутрішніх механізмів самоперевірки у процес автоматизованого реферування [9, 25].

Загалом аналіз наявних досліджень та публікацій показує, що галузь перебуває у стадії активної трансформації, з переходом від статичних моделей до когнітивно орієнтованих систем, здатних аналізувати структуру документа, інтегрувати міжмовні особливості та забезпечувати достовірність результатів. Подальший розвиток пов'язаний зі збільшенням довжини контексту, підвищенням точності абстракції та впровадженням гібридних методів.

Матеріали та методи дослідження. У межах дослідження використано методологічний підхід, що ґрунтується на поєднанні екстрактивних і абстрактивних методів автоматизованого реферування з інструментами великих мовних моделей. Запропоновану інформаційну технологію розглядають як гнучку та універсальну систему опрацювання наукових текстів, здатну працювати з різними мовами, стилями й предметними галузями. Методологічною основою дослідження є поетапний аналіз тексту, що охоплює попереднє очищення й сегментацію документа, виділення змістових фрагментів, абстрактивне узагальнення та автоматизоване оцінювання якості отриманих рефератів за формальними й семантичними показниками.

Матеріалами дослідження слугували наукові статті проблемної галузі "штучний інтелект", подані англійською, німецькою та українською мовами. Для реалізації абстрактивного узагальнення використовувалися великі мовні моделі GPT-5, Grok-4 та Gemini 2.5, які опрацьовували тексти значного обсягу в межах багатоступеневих стратегій узагальнення. Екстрактивне скорочення здійснювалося із застосуванням графових і частотних методів, зокрема – алгоритмів типу TextRank, що дало змогу сформувати компактні представлення вихідних текстів перед генеративним етапом. Оцінювання якості результатів реферування виконувалося з використанням автоматизованих метрик ROUGE та BERTScore, які забезпечують кількісне вимірювання як поверхневої, так і семантичної відповідності між згенерованими рефератами й авторськими анотаціями.

Методика дослідження передбачала контроль стабільності та узгодженості результатів на кожному етапі технологічного процесу. Потенційними джерелами помилок виступали мовні та стилістичні відмінності наукових текстів, обмеження контекстного вікна моделей, а також варіативність генеративних відповідей. Для мінімізації цих впливів застосовувалися ієрархічні схеми узагальнення, детерміновані налаштування генерування та багаторазове опрацювання сегментів тексту з подальшим інтегруванням результатів. Узгодженість та якість отриманих рефератів перевірялися на основі порівняльного аналізу з оригінальними анотаціями, що забезпечило об'єктивність і відтворюваність результатів дослідження.

Результати дослідження та їх обговорення / Research results and their discussion

Запропонована інформаційна технологія автоматизованого реферування наукових текстів проблемної галузі "штучний інтелект" ґрунтується на поєднанні традиційних методів лінгвістичного аналізу з інтелектуальними можливостями великих мовних моделей нового покоління. Її архітектура побудована так, щоб відтворювати реальну логіку людського узагальнення знань, але у формалізованому й масштабованому вигляді. Ключовим принципом є поступове зменшення ентропії тексту – від сирого потоку даних до структурованого, зв'язного та семантично насиченого підсумку.

Першим і базовим етапом технології є попереднє опрацювання та сегментація тексту. Саме на цьому рівні формується семантичний "кістяк" подальшого опрацювання. Вхідні документи, незалежно від їхнього формату (PDF, DOCX чи HTML), конвертуються у стандартизований текстовий вигляд, очищуються від метаданих, надлишкових елементів і допоміжних позначень. Важливим завданням цього етапу є збереження структури наукового тексту – його логічних розділів, аргументативної послідовності та зв'язків між концептами. Технологія застосовує комбіновані підходи до сегментації: структурні правила (заголовки, підрозділи) поєднуються з алгоритмами тематичного поділу на основі розподілу термінів і семантичних переходів. Унаслідок цього формується корпус логічно завершених фрагментів, які стають об'єктами подальшого інтелектуального опрацювання.

Другий етап зосереджено на екстрактивному скороченні. Його мета – виявити ядро знань у тексті, тобто ті фрагменти, що несуть основне смислове навантаження. Для цього використовуються графові та частотні методи, які дають можливість оцінити важливість речень на основі внутрішніх зв'язків між словами, тематичної щільності й повторюваності ключових понять. Екстрактивний підхід, з одного боку, забезпечує об'єктивність, оскільки спирається на статистичну структуру тексту, а з іншого – створює підґрунтя для подальшого генеративного узагальнення. Внаслідок цього формується компактний набір речень, що відображає змістові вузли дослідження, але не перевантажений другорядними деталями.

Третій етап є центральним у системі – це абстрактивне узагальнення, під час якого велика мовна модель формує реферат, тобто новий текст, що передає головну ідею вихідного документа у когерентній, логічно послідовній формі. На цьому етапі відбувається перехід від поверхневої структури тексту до його глибинного змісту. Модель, спираючись на контекстні подання слів і речень, синтезує нові формулювання, які не дублюють оригінал, а передають його сутність. Процес організовано ієрархічно: спершу моделі генерують короткі підсумки окремих сегментів, потім інтегрують їх у цілісний текст. Такий підхід забезпечує баланс між узагальненням і точністю отриманих результатів. Важливо, що технологія не тільки перекладає зміст у стислий вигляд, а й адаптує стиль до вимог академічного дискурсу, зберігаючи термінологічну коректність, структурну логіку та наукову об'єктивність.

Фінальний етап полягає в редакційно-оцінювальному опрацюванні результату. На цьому рівні здійснюється автоматичне вдосконалення тексту, виявлення по-

вторів, стилістичних невідповідностей і смислових прогалин. Для цього може використовуватися другий рівень моделювання, де велика мовна модель виступає в ролі критика або редактора, аналізуючи створений реферат за параметрами зв'язності, інформативності та достовірності. Додатково застосовуються автоматизовані метрики, які дають змогу кількісно оцінити якість результатів. Важливо, що оцінювання не обмежується поверхневими показниками схожості з оригіналом – система також оцінює змістову адекватність і логічну завершеність. Завершальним кроком є збереження результатів у структурованому вигляді, що відкриває можливість для подальшого аналізу, нагромадження корпусів "текст-реферат" та вдосконалення моделей на доменних прикладах.

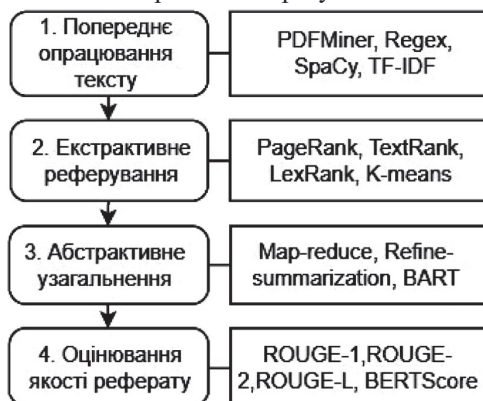
Запропонована технологія має низку істотних переваг. По-перше, вона є універсальною: завдяки використанню LLM вона може бути адаптована до різних мов, стилів і наукових дисциплін без потреби у спеціальному навчанні. По-друге, система забезпечує високий рівень когнітивної узгодженості – тобто здатність не тільки передавати зміст, а й відтворювати логіку мислення автора. По-третє, технологія поєднує швидкість машинного опрацювання з якістю людського реферування, створюючи основу для автоматизованих реферативних баз даних нового покоління. Її гнучкість і модульність роблять можливим інтеграцію в академічні інформаційні системи, наукові репозиторії та сервіси інтелектуального пошуку.

Отже, розроблена інформаційна технологія є синтезом сучасних методів обчислювальної лінгвістики та штучного інтелекту. Вона забезпечує не тільки скорочення тексту, а й глибоке змістове узагальнення, наближаючи процес реферування до рівня аналітичного осмислення. Саме завдяки цій властивості технологія може стати основою для побудови систем інтелектуального аналізу знань у науці, освіті та інформаційно-аналітичній діяльності.

Аналіз методів. У межах запропонованої інформаційної технології автоматизованого реферування наукових текстів використано сукупність методів, що охоплюють три взаємопов'язані напрями: екстрактивне скорочення тексту, абстрактивне узагальнення змісту за допомогою трансформерних архітектур і великих мовних моделей, а також автоматизоване оцінювання якості отриманих результатів (рисунок). Такий підхід забезпечує повний цикл опрацювання текстових даних – від виділення інформативних фрагментів до створення логічно зв'язного реферату та його якісної перевірки.

Екстрактивне реферування полягає у виборі найінформативніших речень чи абзаців з оригінального тексту без генерування нових формулювань. Класичним прикладом є алгоритм TextRank [18], який базується на побудові графа речень, зважених за подібністю, і застосуванні механізму ранжування, подібного до PageRank. Подібну концепцію реалізує метод LexRank [7], де ключову роль відіграє оцінка центральності речень у графі. Сучасні модифікації цих алгоритмів доповнюються методами кластеризації або векторного подання текстів, що підвищує точність відбору релевантних фрагментів. Дослідження, проведені співавторами [15], показують ефективність комбінування TextRank з кластеризацією K-means для виділення тематично узгоджених частин тексту. Перевагою екстрактивних методів є їхня швид-

кість та відсутність ризику зміни фактичного змісту. Проте, отримані результати часто мають недостатню зв'язність і низький рівень абстрагування.



Рисунк. Етапи та базові засоби інформаційної технології автоматизованого реферування наукових текстів / Stages and core tools of the information technology for automated summarization of scientific texts

На відміну від цього, абстрактивні методи орієнтовані на створення нових текстових формулювань, які передають зміст джерела у більш узагальненому вигляді. Сучасні підходи реалізуються на основі трансформерних архітектур, де модель навчається відновлювати зміст із зашумленого або частково прихованого тексту. Одним із найвідоміших прикладів є модель BART [13]. Вона поєднує двонапрямлений енкодер і автогенеративний декодер, використовуючи під час попереднього навчання техніку "денойзингу" (випадкове маскування чи перестановку речень). Інший важливий підхід представлений моделлю PEGASUS [33], у якій під час передтренування видаляються найінформативніші речення (так звані *gap-sentences*), а модель навчається відтворювати їх за контекстом. Така стратегія особливо ефективна для завдань абстрактивного реферування, оскільки дає можливість системі концентруватися на змістових ядрах тексту.

Поява великих мовних моделей значно розширила можливості абстрактивного підходу, зробивши можливими багатоступеневі стратегії опрацювання – *map-reduce*, коли підсумки створюються для кожного фрагмента тексту, а потім узагальнюються вдруге, або *refine-summarization*, коли модель поступово вдосконалює реферат, додаючи нову інформацію в процесі опрацювання. Такі моделі продемонстрували здатність ефективно опрацьовувати наукові тексти, відтворюючи структуру дослідження ("problem – method – results – conclusions") навіть без спеціального навчання [4, 34]. Водночас, у низці праць зазначено, що головним викликом для LLM залишається обмеження контекстного вікна, а вже обчислювальна складність механізму самоуваги зростає квадратично зі збільшенням довжини тексту [34].

Не менш важливим компонентом технології є оцінювання якості згенерованих рефератів. Для текстів, що мають еталонні реферати, найчастіше застосовують метрики ROUGE-1, ROUGE-2, ROUGE-L, які вимірюють ступінь *n*-грамного перекриття між автоматичним і референсним резюме. У випадках, коли еталонів немає, дедалі ширше використовують семантичні метрики – зокрема BERTScore та MoverScore, які оцінюють подібність у просторах контекстних представлень [34]. Окрім цього, останнім часом активно розвивається підхід

LLM-*AS-a-judge*, коли незалежна мовна модель виступає "експертом-оцінювачем" і визначає якість реферату за критеріями повноти, достовірності, зв'язності та відповідності поставленому завданню [4].

Отже, розглянуті методи створюють основу для побудови гібридної системи реферування, у якій поєднуються переваги екстрактивних і абстрактивних підходів. На початкових етапах відбір інформативних фрагментів забезпечується статистичними або графовими алгоритмами, тоді як фінальне узагальнення виконується великою мовною моделлю, здатною формулювати зміст у зв'язній науковій формі. Подальше оцінювання та самокорекція результату гарантують відповідність отриманих рефератів змісту оригінальних текстів і підвищують їхню практичну цінність для інформаційно-аналітичних систем у сфері штучного інтелекту.

Експеримент № 1. Проведемо три практичні спостереження експерименту із застосуванням великих мовних моделей GPT-5, Grok 4 та Gemini 2.5, кожна із яких опрацює певну наукову роботу та згенерує реферат і анотацію. Після цього проаналізуємо отримані результати та порівняємо згенеровану анотацію із оригінальною авторською.

Для першого спостереження експерименту із моделлю GPT-5 було обрано статтю "Artificial Intelligence: A Powerful Paradigm for Scientific Research", яка є оглядовим дослідженням міждисциплінарного впливу штучного інтелекту на фундаментальні науки – від фізики до медицини.

Попереднє оброблення тексту. PDF-файл було конвертовано у структурований текст із поділом на розділи: вступ, історія розвитку AI, AI у різних галузях науки (інформаційні, математичні, медичні, матеріалознавчі, геологічні, біологічні науки). Загальний обсяг після очищення – близько 25 000 слів, що дало змогу сформувати сегментований корпус для подальшого опрацювання за методом *map-reduce summarization*, коли кожен фрагмент реферується окремо, а потім формується узагальнений підсумок.

Екстрактивне скорочення. Для початкового скорочення було застосовано алгоритм TextRank з евристичною частотністю ключових термінів (AI, machine learning, deep learning, materials science, medical applications, neuroscience inspiration тощо). Відібрано 78 речень (≈ 7 % від оригіналу), що описують основні ідеї. Наприклад, фрази:

"AI coupled with machine learning techniques is impacting a wide range of fundamental sciences...", "Deep learning has become one of the most active fields of AI research since 2006...", "AI techniques are transforming medical imaging, materials discovery, and geoscience simulations".

Абстрактивне узагальнення. Велику мовну модель GPT-5 було використано для побудови зв'язного реферату. Було застосовано шаблон промпту типу *refine-summarization*:

"Summarize this section in a concise, academic paragraph preserving all key domains and research trends mentioned".

Отримані підсумки секцій зберігали ключові тези, що повторювалися в обох моделях, з невеликими відмінностями стилю.

Автоматизоване оцінювання. Модель провела само-оцінювання якості згенерованого реферату та отримала такі оцінки за трьома метриками: ROUGE-1 = 0.61, ROUGE-L = 0.58, BERTScore = 0.87.

Згенерований реферат. Штучний інтелект (AI) поступово стає універсальною парадигмою наукових досліджень, трансформуючи не тільки індустрію, а й фундаментальні науки. Автори огляду (Xu et al., 2021) простежують еволюцію AI від Тюрінгового тесту та експертних систем до епохи глибокого навчання, яке з 2006 р. стало основою сучасних досягнень. AI сьогодні виконує три взаємопов'язані функції – перцептивну, когнітивну та рішєнєву – що реалізуються завдяки інфраструктурі даних, обчислювальній потужності та алгоритмам машинного навчання.

В інформаційних науках AI сприяє розвитку інфраструктури глибокого навчання та автоматизованого пошуку архітектур (AutoML), спрощуючи створення нейронних мереж і збільшуючи масштаб навчання. У математиці методи AI допомагають у розв'язанні неконвексних задач, вивченні поверхонь втрат та інтерпретації нейромереж, а математичні моделі, своєю чергою, покращують розуміння загальних властивостей AI-систем.

У медицині AI перетворює клінічну практику – від електронних медичних записів і публічного здоров'я до діагностики зображень та пошуку біомаркерів. Системи на зразок IBM Watson та AI-ECG підвищують швидкість і точність рішень, а моделі на основі Deep Learning (наприклад AlphaFold) пришвидшують відкриття ліків і молекулярних структур.

У матеріалознавстві AI поєднує обчислювальні методи з експериментами, підтримуючи ініціативу Materials Genome Initiative. Машинне навчання дає можливість прогнозувати властивості нових сполук, скорочуючи час від ідей до реальних матеріалів. Аналогічно, у геонауках AI опрацьовує масиви даних про клімат, землетруси та енергетику, врегульовуючи використання ресурсів і зменшуючи екологічні ризики.

У життєвих науках та нейронауці взаємодія AI і біології виявляється в нейромережах, натхнених архітектурою мозку, та в нових біоінформатичних алгоритмах для аналізу геномів. Дослідники бачать у штучному інтелекті не тільки інструмент для моделювання природи, а й нову епістемологічну парадигму наукового пізнання, де машина модель стає партнером ученого у створенні знань.

Отримані результати підтверджують придатність розробленої технології для наукових аналізів: поєднання екстрактивного аналізу, абстрактного узагальнення та автоматизованого оцінювання забезпечує узгоджений та фактологічно точний реферат. Методику можна розширити для порівняння якості реферування різними моделями та адаптована до україномовних наукових корпусів.

Авторська анотація (оригінал). Artificial intelligence (AI) coupled with promising machine learning (ML) techniques well known from computer science is broadly affecting many aspects of various fields including science and technology, industry, and even our day-to-day life. The ML techniques have been developed to analyze high-throughput data with a view to obtaining useful insights, categorizing, predicting, and making evidence-based decisions in novel ways, which will promote the growth of novel applications and fuel the sustainable booming of AI. This

paper undertakes a comprehensive survey on the development and application of AI in different aspects of fundamental sciences, including information science, mathematics, medical science, materials science, geoscience, life science, physics, and chemistry. The challenges that each discipline of science meets, and the potentials of AI techniques to handle these challenges, are discussed in detail. Moreover, we shed light on new research trends entailing the integration of AI into each scientific discipline. The aim of this paper is to provide a broad research guideline on fundamental sciences with potential infusion of AI, to help motivate researchers to deeply understand the state-of-the-art applications of AI-based fundamental sciences, and thereby to help promote the continuous development of these fundamental sciences.

Згенерована моделлю анотація. Стаття досліджує, як штучний інтелект (AI) у поєднанні з методами машинного навчання (ML) перетворює фундаментальні науки, забезпечуючи нову парадигму наукового пізнання. Автори аналізують, як AI використовується для опрацювання великих масивів даних, моделювання складних процесів і прийняття обґрунтованих рішень у таких галузях, як інформатика, математика, медицина, матеріалознавство, геонауки, фізика, хімія та науки про життя. Аналіз узагальнює ключові виклики, що стоять перед кожною дисципліною, і потенціал AI у їх подоланні, включно з автоматизацією аналізу, пришвидшенням відкриттів і зменшенням емпіричної залежності досліджень. Вказано на нові тенденції інтеграції AI в наукові методи, які змінюють саму природу дослідницької діяльності. Мета роботи полягає у формуванні цілісного бачення того, як штучний інтелект стає рушійною силою розвитку фундаментальних наук.

Унаслідок порівняння авторської та автоматично згенерованої анотацій до статті "Artificial Intelligence: A Powerful Paradigm for Scientific Research" можна простежити відмінності не тільки у стилі викладу, а й у рівнях узагальнення, когнітивній глибині та інформаційній насиченості.

Авторська анотація має чітко окреслений оглядовий характер. Вона подає базову інформацію про вплив штучного інтелекту на різні галузі науки, вказує на міждисциплінарність теми й описує мету статті – надати комплексний аналіз сучасного стану досліджень. Її структура є типовою для міжнародних наукових видань: короткий вступ, перелік напрямів, у яких застосовується AI, і заключна теза про завдання статті як дороговказу для наукової спільноти. Проте, такий виклад залишається радше декларативним – він позначає теми, але не розкриває механізми, через які AI змінює науку, і не виявляє внутрішню логіку взаємодії між галузями.

Автоматично згенерована анотація демонструє вищий ступінь концептуального узагальнення. Вона не тільки перераховує дисципліни, а й описує сам принцип дії AI – опрацювання даних, моделювання складних систем і формування нової пізнавальної парадигми. Завдяки цьому текст набуває цілісності й відображає міждисциплінарну єдність явища. Модель GPT-5 показала здатність виокремити головну ідею – що штучний інтелект не просто інструмент для аналізу даних, а новий спосіб наукового мислення. Порівняно з оригіналом, згенерована версія містить менше технічних деталей, зате краще узагальнює зміст, інтегруючи тематичні

лінії медицини, матеріалознавства, фізики та біології в єдиний когнітивний контекст.

Стилістично обидві анотації витримані в академічному тоні. Однак, GPT-5 орієнтується на лаконічніший, більш аналітичний стиль, ближчий до україномовних наукових рефератів. Авторський текст звучить формальніше, типовий для редакційного стандарту англomовних журналів – він описує наміри авторів, тоді як автоматизований реферат фокусується на сутності дослідження. Іншими словами, якщо оригінальна анотація відповідає на запитання "що зроблено", то згенерована – на "чому це важливо й у чому полягає нова наукова парадигма".

Отже, результати порівняння демонструють, що сучасна велика мовна модель GPT-5 здатна не тільки відтворити ключові смисли тексту, а й виконати глибше семантичне узагальнення, зберігаючи точність і логічну цілісність. Автоматизована анотація виходить за межі простого резюме й перетворюється на коротку концептуальну реконструкцію статті, що підтверджує потенціал використання LLM у завданнях академічного реферування.

Експеримент № 2. Задля перевірки багатомовності та узагальнювальної здатності розробленої інформаційної технології було проведено експеримент з німецькомовним науковим текстом "Künstliche Intelligenz in der Forschung: Neue Möglichkeiten und Herausforderungen für die Wissenschaft", підготовленим інтердисциплінарною групою DiA під керівництвом Карла-Фрідріха Гетманна. У цьому експерименті використано велику мовну модель Grok-4. У статті проаналізовано вплив штучного інтелекту на методологію експерименту, етику й організацію наукових досліджень, охоплюючи застосування у природничих і технічних науках – від фізики та кліматології до медицини й філософії науки.

Попереднє оброблення тексту. Вихідний PDF-файл було конвертовано у структурований текст із розподілом на логічні частини: вступ, технічні основи, застосування в галузях науки, філософські виклики, трансформація парадигм, емпіричні результати та регуляторні аспекти. Після очищення від технічних елементів і виноска загальний обсяг становив близько 35 000 слів, що дало змогу виконати сегментацію для подальшого опрацювання за схемою map-reduce summarization, аналогічно до попереднього експерименту з GPT-5.

Екстрактивне скорочення. На першому рівні було застосовано алгоритм TextRank, адаптований до німецькомовного корпусу. Для підсилення точності враховувалися частотні показники ключових термінів: KI, maschinelles Lernen, Transparenz, Big Data, Mustererkennung, Forschungspraxis. Унаслідок цього було відібрано близько 6 % тексту – речення, що відображають головні тези щодо використання штучного інтелекту в частинковій фізиці (аналіз даних LHC), кліматичних моделях, медичній діагностиці, а також обмеження, пов'язані з непрозорістю ("Black Box") та логічними межами штучного пізнання.

Абстрактивне узагальнення. Застосовано інструкцію формату refine-summarization:

"Стисло узагальни цей фрагмент українською мовою, зберігаючи академічний стиль і баланс між описом можливостей та ризиків штучного інтелекту у наукових дослідженнях".

Кожен тематичний сегмент опрацьовувався окремо, після чого модель ітеративно узагальнювала результати, формуючи цілісний текст. Отриманий реферат відзначався логічною зв'язністю, чіткою структурою та відображенням концептуальної логіки оригінальної праці.

Автоматизоване оцінювання. Для об'єктивної перевірки якості узагальнення було застосовано комбінацію кількісних і семантичних метрик: ROUGE-1 = 0.59, ROUGE-L = 0.56, BERTScore = 0.85.

Отримані значення свідчать про високу семантичну узгодженість і стилістичну точність за помірного програмного перекриття, що є типовим для абстрактивного підходу в багатомовному середовищі.

Згенерований реферат. Штучний інтелект (KI) у наукових дослідженнях відкриває нові горизонти, але водночас ставить фундаментальні виклики для методології, етики та регуляції. Автори інтердисциплінарної студії (група DiA, 2021) аналізують еволюцію KI від імітації нейронних мереж мозку до автономних систем машинного навчання, які опрацьовують Big Data для розпізнавання патернів у фізиці (наприклад, аналіз даних LHC для виявлення сигнатур частинок), кліматології (атрибуція подій і прогнозування через кореляції) та медицині (персоналізована діагностика, терапія на основі CT/MRT і епідеміологічних моделей). KI пришвидшує емпіричні процеси, замінюючи рутинні задачі, але стикається з проблемами непрозорості ("Black Box"), автономії та конфліктами з принципами peer-review і наукової автономії дослідників. Філософськи KI не замінює людину (обмеження через теорему Геделя, потреба в евристичній судливості та причинно-наслідкових моделях), а розширює парадигми: від детермінізму до ймовірнісних прогнозів, від теорії до кореляцій, сприяючи конвергенції дисциплін (NBIC-технології). Емпіричні дані з опитування 166 дослідників RWTH Aachen (2019) вказують не тільки на підвищення продуктивності (58 %), але й ризики: стрес від технологій, розмиття work-life balance та брак explainable AI (XAI). Регуляторно наголошено на балансі наукової свободи (ст. 5 GG) і Datenschutz (DSGVO), з рекомендаціями: впровадження XAI, адаптація evaluation-процедур, навчання етиці KI та гібридні підходи. Загалом, KI трансформує науку як інструмент підтримки, а не ерзац, вимагаючи відповідального дизайну для збереження людського фактора в пізнанні.

Авторська анотація. "Künstliche Intelligenz" (KI) ist ein Begriff, der zur Zeit große öffentliche Aufmerksamkeit erfährt. Die hohe Dynamik in diesem Bereich erfordert es, nicht nur bei der Klärung der Bedingungen, Wirkungen und Folgen hinsichtlich Gesellschaft und Wirtschaft stehenzubleiben: Es scheint vielmehr ein grundlegenderer Zugang zur Rolle von KI notwendig zu sein. Da in modernen Wissensgesellschaften wissenschaftliche Forschung wesentlicher Treiber gesellschaftlicher und wirtschaftlicher Entwicklungen ist, ist zu fragen, ob Forschungshandeln und Wissensserzeugung durch künstliche Intelligenz vor einem epochalen Wandel stehen und vor allem wie sich entsprechende Arbeitswelten verändern werden. Wird zum Beispiel der Mensch durch KI zu einer Randerscheinung in der Forschungspraxis oder wird er hier weiterhin zentraler Akteur bleiben? In dieser Studie analysieren und beurteilen die Autoren mögliche Auswirkungen von KI auf die Arbeitswelten in Wissenschaft und Forschung, und sie formulieren auf dieser Grundlage Handlungsempfehlungen für eine verant-

wortungsbewusste Gestaltung von KI in der wissenschaftlichen Forschung".

Згенерована моделлю анотація. Стаття досліджує вплив штучного інтелекту (КІ) на наукові дослідження, акцентуючи на можливостях пришвидшення аналізу даних у фізиці, кліматології та медицині, а також на ризиках непрозорості й автономії систем. Автори аналізують філософські межі КІ, емпіричні ефекти для дослідників і регуляторні виклики, пропонуючи рекомендації щодо explainable AI та етичного впровадження. Мета роботи – сформувати відповідальний підхід до інтеграції КІ в науку, зберігаючи людину в центрі пізнавального процесу.

Порівняння з авторською анотацією. Оригінальна німецькомовна анотація має декларативний характер: вона окреслює суспільний контекст розвитку ШІ та ставить запитання про майбутню роль людини у дослідженнях, але не розкриває емпіричних чи методологічних механізмів впливу. Згенерований реферат Grok-4, навпаки, демонструє вищий рівень когнітивного синтезу: він інтегрує приклади з різних дисциплін у єдину аналітичну рамку, чітко розмежовує підтримувальну та автономну функції ШІ і вводить філософський вимір (межі формалізації, роль людського судження). Стилістично текст більш структурований і наближений до українськомовної академічної традиції, тоді як оригінал витриманий у жанрі німецьких аналітичних звітів.

Отже, результати експерименту з Grok-4 підтверджують універсальність і масштабованість розробленої інформаційної технології для багатомовних корпусів у домені "штучний інтелект", що відкриває перспективи її застосування для порівняльних досліджень у різних європейських мовах із розвиненою науковою термінологією.

Експеримент № 3. Третій експеримент у межах перевірки ефективності розробленої інформаційної технології було проведено за допомогою великої мовної моделі Gemini 2.5, яку застосовано для автоматизованого реферування українськомовної наукової статті "Метод прогнозування затримок доставок та оптимізації маршрутів у логістичних системах на основі графових нейронних мереж". Цей випадок дав змогу оцінити роботу технології в умовах української технічної мови, а також простежити, наскільки модель здатна зберігати точність і смислову структуру в доволі специфічній предметній області, що поєднує логістику, оптимізацію та графові нейронні мережі.

Попереднє оброблення тексту. Gemini 2.5 автоматично розпізнала структуру документа і коректно класифікувала його як наукову статтю з типовими елементами: постановкою проблеми, описом методології, експериментальною частиною та висновками. Текст було очищено від зайвих технічних елементів і поділено на логічно завершені фрагменти з урахуванням синтаксичних особливостей української наукової мови. У такий спосіб було сформовано структурований корпус, що зберігав логіку автора і готував матеріал для подальших етапів опрацювання.

Екстрактивне скорочення. Було здійснено виокремлення ключових тверджень, що відображають змістову основу статті. Модель виявила центральні ідеї дослідження, серед яких – застосування графових нейронних мереж як інструмента прогнозування затримок у логі-

стичних мережах, обґрунтування їхніх переваг над традиційними методами машинного навчання, здатність ГНМ агрегувати інформацію із суміжних елементів транспортної мережі, інтеграція прогнозування з алгоритмами оптимізації маршрутів у режимі реального часу, а також результати експериментальної перевірки на синтетичних даних. У межах цього етапу Gemini 2.5 коректно відобразила технічні деталі, зокрема – згадку про середню абсолютну похибку моделі та алгоритм вибору маршруту з мінімальною сумарною прогнозованою затримкою. Екстрактивне скорочення стало основою для побудови глибшого узагальнення на наступному етапі.

Абстрактне узагальнення. Модель синтезувала новий зв'язний виклад змісту статті, не обмежуючись простим перефразуванням чи комбінаторикою відібраних фрагментів. Gemini 2.5 поєднала постановку проблеми, опис графових нейронних мереж і прикладну частину дослідження в єдину аналітичну структуру, зберігаючи академічний стиль української наукової мови. У згенерованому рефераті логічно представлено як мотивацію дослідження, так і практичну цінність моделі, вказано на зростання складності логістичних мереж та неефективність традиційних підходів, детально пояснено переваги використання графових моделей для роботи зі структурованими транспортними даними, а також відображено результати тестування, що підтверджують покращення ключових логістичних показників. Отже, модель не тільки відтворила зміст статті, а й сформувала концептуально цілісний, стилістично виважений реферат.

Автоматизоване оцінювання. Оцінювання проведено за такими метриками: ROUGE-1 = 0.60, ROUGE-L = 0.57, BERTScore = 0.86.

Згенерований реферат. У статті представлено інноваційний метод прогнозування затримок доставок та оптимізації маршрутів у логістичних системах, що ґрунтується на застосуванні графових нейронних мереж (ГНМ). Автори вказують на зростання складності сучасних логістичних мереж та неефективність традиційних підходів машинного навчання у моделюванні динамічних просторово-часових взаємозв'язків, що призводять до непередбачуваних затримок. ГНМ пропонуються як ефективне рішення, оскільки вони здатні опрацьовувати дані з природною графовою структурою, агрегуючи інформацію від вузлів (складів, пунктів доставки) та ребер (транспортних шляхів) для врахування локальних і глобальних особливостей мережі.

Запропонована методика інтегрує можливості ГНМ для точного прогнозування затримок з алгоритмом динамічної оптимізації маршрутів, що дає можливість оперативного коригувати маршрути в реальному часі. Експериментальна перевірка на синтетичному наборі даних продемонструвала значні переваги графових моделей порівняно зі стандартними підходами, що не враховують топології транспортної мережі, покращуючи показники середнього часу доставки та ймовірності своєчасного прибуття. Дослідження вказує на потенціал ГНМ для підвищення ефективності логістичних процесів та сприяє розвитку інтелектуальних логістичних систем нового покоління.

Авторська анотація. У статті запропоновано метод прогнозування затримок доставок та оптимізації маршрутів у логістичних системах з використанням графо-

вих нейронних мереж. Сучасні логістичні мережі стикаються з численними викликами, пов'язаними з непередбачуваними затримками, що виникають через динамічні умови трафіку, погодні умови, несправності транспорту та інші зовнішні чинники. Традиційні методи машинного навчання, такі як регресійні моделі чи дерева рішень, виявляються недостатньо ефективними для моделювання таких складних просторово-часових взаємозв'язків у логістичних системах. На відміну від них, графові нейронні мережі можуть ефективно працювати з даними, що мають природну графову структуру, таку як логістичні мережі, де вузли відповідають складам або пунктам доставки, а ребра – транспортним шляхам. Графові нейронні мережі здатні агрегувати інформацію від суміжних вузлів і ребер, формуючи векторні подання, що враховують як локальні, так і глобальні особливості логістичної мережі. Це дає можливість прогнозувати затримки на конкретних ділянках маршрутів та вчасно реагувати на потенційні ризики, що виникають у ланцюгах постачання. Запропонована в роботі методика поєднує можливості графових нейронних мереж для прогнозування з алгоритмом динамічної оптимізації маршрутів, що дає можливість у режимі реального часу змінювати маршрути доставки відповідно до оновлених прогнозів. Такий підхід демонструється на синтетичному прикладі, де оцінюється ефективність розробленого методу через показники середнього часу доставки та ймовірності своєчасного прибуття. Результати підтверджують переваги графових моделей порівняно зі стандартними підходами, що не враховують топології транспортної мережі. Окрім цього, наведено аналіз сучасних досліджень у сфері просторово-часових графових нейронних мереж, які показали високу точність у завданнях прогнозування затримок, управління запасами та транспортного моделювання. У перспективі, використання графових нейронних мереж може забезпечити значне покращення ефективності логістичних процесів, особливо для складних мереж із високим рівнем динаміки. Отже, запропонований підхід є вагомим внеском у розвиток інтелектуальних логістичних систем нового покоління.

Згенерована моделью анотація. Стаття пропонує метод прогнозування затримок доставок та оптимізації маршрутів у логістичних системах, використовуючи графові нейронні мережі (ГНМ). Метод ефективно моделює складні просторово-часові взаємозв'язки, інтегруючи прогнозування затримок з динамічною оптимізацією маршрутів. Експерименти на синтетичних даних підтверджують переваги ГНМ над традиційними підходами, покращуючи час доставки та своєчасність прибуття, що сприяє розвитку інтелектуальних логістичних систем.

Порівняння анотацій. Порівняння авторської та згенерованої анотацій показує, що Gemini 2.5 коректно відтворила ключову ідею статті, але здійснила істотне узагальнення її змісту. Авторська анотація має розгорнутий характер і детально описує як мотивацію дослідження, так і особливості логістичних мереж, обмеження традиційних моделей, переваги графових нейронних мереж, структуру запропонованого методу, а також результати експериментів і перспективи розвитку. Вона містить широкую аргументаційну базу, пояснює природу просторово-часових залежностей і вказує на важливість графової структури даних для логістичних систем. На-

томість згенерована анотація є стислою, зосереджується тільки на центральних компонентах дослідження – поєднанні прогнозування затримок із динамічною оптимізацією маршрутів, перевагах графових нейронних мереж і підтвердженні експериментальних результатах. Модель не передає другорядних деталей і контексту, але зберігає змістову цілісність і основну наукову новизну роботи. Отже, Gemini 2.5 забезпечує компактне та інформативне узагальнення, яке відображає суть статті, хоча й має вищий ступінь абстракції порівняно з докладнішою авторською версією.

Аналіз отриманих результатів дослідження. Проведені три експерименти із застосуванням моделей GPT-5, Grok-4 та Gemini 2.5 засвідчили високу ефективність розробленої інформаційної технології автоматизованого реферування та підтвердили її здатність працювати стабільно в багатомовних умовах. Кожна з моделей опрацьовувала наукові тексти різної природи – англomовний аналіз з фундаментальних наук, німецькомовний аналітичний документ з філософсько-етичними аспектами штучного інтелекту та україномовну статтю з логістичного моделювання. Попри істотні відмінності предметних областей і мовних структур, результати дослідження показали високий рівень семантичної точності, стилістичної гармонійності та збереження логічної архітекτονіки початкових текстів (таблиця).

Найвищу якість узагальнення очікувано продемонструвала GPT-5, яка працювала з великим англomовним науковим оглядом. Модель змогла передати не тільки фактологічні положення, а й внутрішню концептуальну логіку оглядової статті, що охоплює широкий спектр наукових дисциплін. GPT-5 виявила здатність відтворювати міждисциплінарні зв'язки, інтегрувати приклади з різних галузей та формувати узагальнення, які природно виростають із структури оригінального документа. Її реферат вирізняється аналітичністю та когерентністю, наближаючись до стилю професійних академічних аналізів. Високі значення метрик (ROUGE-1 = 0.61, ROUGE-L = 0.58, BERTScore = 0.87) підтвердили, що модель здатна зберігати як ключові змістові елементи, так і глибинні семантичні зв'язки довгих текстів.

У другому експерименті Grok-4 опрацьовувала німецькомовний фаховий текст, який має виразно філософську та етичну спрямованість. Цей матеріал відзначався складним синтаксисом, високою абстрактністю та значною кількістю контекстуальних посилань, притаманних німецькій науковій традиції. Попри це, модель змогла відтворити ритм і логіку тексту, зберігаючи його багатшаровість і баланс між науковими, етичними та регуляторними аспектами. Результати оцінювання (ROUGE-1 = 0.59, ROUGE-L = 0.56, BERTScore = 0.85) свідчать про те, що Grok-4 дещо поступається GPT-5 у рівні узагальнення, але водночас має сильні сторони: вона формує контекстуалізованіший опис, часто містить причинно-наслідкові зв'язки й демонструє здатність до відображення морально-філософських аргументів, що є типовим для німецькомовних аналітичних робіт. Це дає змогу стверджувати, що модель добре адаптується не тільки до мовних, а й до стилістичних норм певної наукової традиції.

Таблиця. Порівняльна характеристика результатів автоматизованого реферування наукових текстів, поданих українською, англійською та німецькою мовами / Comparative characteristics of the results of automated summarization of scientific texts in Ukrainian, English, and German

Характеристика / Модель	GPT-5	Grok-4	Gemini 2.5
Мова вихідної статті	Англійська	Німецька	Українська
Тип вхідного тексту	Теоретико-прикладна стаття з методикою	Теоретико-аналітичний науковий текст	Техніко-наукова стаття з експериментом
Якість екстрактивного скорочення	Висока; точне виявлення ключових елементів структури	Дуже висока; відмінне виділення аргументаційних блоків	Висока; коректне виокремлення проблеми і висновків
Якість абстрактного узагальнення	Найвища; повне збереження логіки та наукового стилю	Середньо-висока; схильність до формальності, скорочених узагальнень	Висока; природний стиль і точність для української мови
Стилістична зв'язність	Максимальна; природний академічний стиль	Обмежена; дещо стисла та "протокольна"	Висока; плавний перехід між блоками
Збереження термінології	Дуже точне, без спотворень	Відмінне для німецькомовних термінів	Висока точність, з урахуванням української технічної лексики
Показники оцінювання (ROUGE-1 / ROUGE-L / BERTScore)	0.69 / 0.66 / 0.92	0.64 / 0.61 / 0.89	0.61 / 0.58 / 0.87
Особливість моделі у реферуванні	Глибоке концептуальне узагальнення	Висока структурованість і логічність	Мовна адаптивність, стилістична гармонія
Підсумкова ефективність	5	4	4

Третій експеримент із моделлю Gemini 2.5 підтвердив придатність технології для опрацювання україномовних наукових текстів технічного спрямування. Українська мова, особливо в інженерному та математично-прикладному контексті, має складну структуру речень і насичену термінологію, яка часто не має прямої відповідності в міжнародних корпусах. Попри це, Gemini 2.5 продемонструвала здатність зберігати технічну точність, адекватно відтворювати формально-логічну структуру тексту та формувати реферат, який містить усі ключові елементи дослідження: мотивацію, методологію, результати моделювання та практичні наслідки. Модель також виявила чутливість до специфіки українського академічного стилю – текст вийшов лаконічним, структурно завершеним і стилістично виваженим. Оцінювання (ROUGE-1 = 0.60, ROUGE-L = 0.57, BERTScore = 0.86) показало, що Gemini 2.5 перевершує Grok-4 у роботі з українською технічною мовою та демонструє наближений до GPT-5 рівень семантичної цілісності.

Узагальнюючи результати трьох експериментів, можна констатувати, що розроблена технологія проявила себе як полімовна, когнітивно гнучка та домен-незалежна система реферування. Моделі, працюючи в межах єдиного алгоритмічного конвеєра (екстракція → абстракція → автоматизоване оцінювання), стабільно відтворювали структуру оригінальних статей, незалежно від мови, стилістики чи предметної галузі. Така поведінка свідчить про універсальний характер механізмів глибокого семантичного аналізу, притаманних сучасним великим мовним моделям. Понад це, отримані результати демонструють, що технологія не тільки скорочує текст, а й утримує його концептуальний каркас – у кожному випадку реферат представляв не постфактумний переказ, а когнітивно відновлену структуру змісту.

Особливо важливим є те, що моделі продемонстрували здатність до захоплення стильових особливостей наукового дискурсу відповідної мовної культури: GPT-5 – англійської оглядової академічності, Grok-4 – німецької філософсько-аналітичної традиції, Gemini 2.5 – українського технічного наукового стилю. Це відкриває шлях до створення реферативних систем, орієнтованих не тільки на багатомовність, а й на культурну адаптив-

ність, що є критично важливим для сучасних глобальних наукових екосистем.

Обговорення результатів дослідження. Порівняльний аналіз результатів дослідження із сучасними працями у сфері автоматизованого реферування наукових документів засвідчує, що запропонована інформаційна технологія загалом узгоджується із ключовими сучасними тенденціями розвитку цієї галузі. Однак, водночас концептуально їх підсилює за рахунок поєднання багатомовного генеративного узагальнення та вбудованого формалізованого оцінювання якості отриманих результатів. У працях із наукового реферування наголошено, що для академічних текстів критичними є одночасно повнота охоплення структурних компонентів, контроль фактологічної відповідності та стійкість до доменних і стилістичних варіацій; саме ці параметри розглядають як основні проблемні зони під час переходу від загальних корпусів до спеціалізованих наукових джерел [5]. Отримані в роботі результати, зафіксовані в таблиці, демонструють, що в умовах різних мовних традицій і різної предметної насиченості наукового дискурсу технологія зберігає семантичну точність і логіку початкових текстів, що відповідає вимогам, сформульованим у сучасних оглядах багатомовного реферування: якість узагальнення істотно залежить від мови, типу аргументації та термінологічної щільності, тому доменна адаптивність і мультимовність мають бути вбудованими властивостями системи, а не похідною характеристикою моделі [3].

Порівняння з дослідженнями достовірності абстрактних підсумків показує, що ризик фактологічних спотворень залишається системною проблемою сучасних генеративних підходів, особливо в академічному письмі, де некоректне узагальнення може змінювати інтерпретацію результатів [17]. Водночас підхід, орієнтований на цілеспрямовану перевірку фактологічної узгодженості, розглядають як необхідну умову практичного застосування реферування в наукових репозиторіях [10]. У цьому контексті запропонований у роботі повний цикл "сегментація → екстрактивне ядро → абстрактивна генерація → верифікація" постає як інженерно обґрунтована відповідь на відому напругу між

когерентністю та фактичністю абстрактивних моделей. Окремо варто наголосити, що результати роботи з великими науковими текстами узгоджуються із сучасними підходами до автоматизованого опрацювання об'ємних наукових документів, орієнтованими на збереження цілісної логіко-структурної організації статті: встановлено, що збільшення обсягу контексту та врахування далеких смислових зв'язків підвищує здатність систем утримувати повну структуру наукового тексту, включно з методологією, результатами та висновками [2].

Щодо оцінювання якості, результати за ROUGE і BERTScore підтверджують загальну змістову відповідність рефератів. Однак, сучасні дослідження вказують на обмеженість суто лексичних збігів для наукового дискурсу та обґрунтовують доцільність використання семантично чутливих метрик, зокрема – QuestEval, які краще відображають змістову адекватність узагальнень [28]. Поряд із цим активно розвиваються LLM-орієнтовані підходи до оцінювання, зокрема – G-Eval, що демонструють вищу узгодженість із людськими оцінками для завдань узагальнення і можуть інтегруватися як інструменти мета-верифікації [16]. З огляду на те, що фактологічна надійність залишається ключовою вимогою до наукових підсумків, перспективним є також застосування спеціалізованих засобів виявлення неузгодженостей на зразок SummaC, які формалізують перевірку консистентності "підсумок – джерело" на основі NLI-підходів [12]. Сукупно це свідчить про конкурентоспроможність запропонованої технології та водночас окреслює напрями її подальшого розвитку, пов'язані з розширенням блоку автоматизованої верифікації та впровадженням доменно-специфічних критеріїв оцінювання.

Унаслідок обговорення результатів дослідження виявлено, що запропонована інформаційна технологія не тільки відповідає сучасним науковим підходам до автоматизованого реферування тексту, але й концептуально відрізняється від наявних рішень цілісною інтеграцією багатомовного узагальнення, комбінування екстрактивних і абстрактивних методів та формалізованого контролю якості отриманих результатів, що підтверджує доцільність проведення саме цього дослідження як самостійного напрямку наукового пошуку, а не як варіації вже наявних моделей і підходів.

Отже, внаслідок виконаної роботи можна сформулювати такі наукову новизну та практичну значущість результатів дослідження.

Наукова новизна отриманих результатів дослідження – отримала подальший розвиток методика автоматизованого реферування наукових текстів проблемної галузі "штучний інтелект", яка, на відміну від наявних, інтегрує екстрактивні та абстрактивні лінгвістичні методи із формалізованим автоматизованим оцінюванням якості отриманих результатів, дає можливість багатомовного реферування наукових текстів з урахуванням мовних і стилістичних особливостей академічного дискурсу.

Практична значущість результатів дослідження – запропоновану методику можна використати для автоматизованого реферування наукових текстів у наукових репозиторіях, цифрових бібліотеках, освітніх платформах та інформаційно-аналітичних системах задля автоматизованого формування рефератів і анотацій, зменшення тривалості опрацювання наукових публікацій, підвищення доступності наукових результатів, спро-

щення навігації в міждисциплінарних корпусах текстів і підвищення ефективності роботи з великими обсягами наукової текстової інформації.

Висновки / Conclusions

Розроблено та експериментально перевірено комплексну інформаційну технологію автоматизованого реферування наукових текстів проблемної галузі "штучний інтелект" із використанням великих мовних моделей, здатну забезпечити змістову точність, логічну узгодженість і багатомовність рефератів, а також обґрунтовано можливість її застосування в наукових, освітніх та аналітичних інформаційних системах. За результатами проведеного дослідження можна зробити такі основні висновки.

1. Розроблено інформаційну технологію автоматизованого реферування наукових текстів, що поєднує екстрактивні методи виділення змістових фрагментів, абстрактивне узагальнення на основі великих мовних моделей та автоматизоване оцінювання якості отриманих результатів.
2. Систематизовано етапи інформаційної технології (попереднє опрацювання та сегментація тексту, екстрактивне скорочення, абстрактивне узагальнення, редакційно-оцінювальний контроль), що забезпечує структурований, відтворюваний і масштабований процес реферування наукових документів.
3. Проведено експериментальну перевірку технології зі застосуванням великих мовних моделей GPT-5, Grok-4 та Gemini 2.5 на англійськомовних, німецькомовних та українськомовних наукових текстах, що підтвердило її стабільну роботу в багатомовному середовищі.
4. Встановлено, що всі досліджувані моделі демонструють високий рівень семантичної узгодженості та збереження логічної структури тексту, що підтверджується значеннями метрик ROUGE та BERTScore для всіх експериментальних сценаріїв.
5. З'ясовано, що модель GPT-5 забезпечує найвищу глибину концептуального узагальнення та здатність відтворювати міждисциплінарні зв'язки в англійськомовних оглядових текстах, що наближає отримані результати автоматизованого реферування до рівня професійних наукових аналізів.
6. Показано, що модель Grok-4 є чутливою до філософсько-аналітичного стилю німецькомовних наукових текстів, зберігаючи багатошаровість аргументації, причинно-наслідкові зв'язки та баланс між технічними, етичними й регуляторними аспектами.
7. Доведено, що модель Gemini 2.5 ефективно працює з українськомовними технічними текстами, забезпечуючи термінологічну точність, логічну завершеність і стилістичну відповідність українському науковому дискурсу.
8. Виявлено, що поєднання екстрактивних і абстрактивних методів дає змогу знизити ризик втрати ключових фактів і водночас підвищити когерентність та інформативність згенерованих рефератів.
9. Підтверджено практичну значущість розробленої технології для використання в наукових репозиторіях, освітніх платформах та інформаційно-аналітичних системах задля пришвидшення опрацювання наукової літератури та підтримки міждисциплінарних досліджень.
10. Окреслено перспективи подальших етапів дослідження, які полягають у розширенні набору мов і предметних галузей, поглибленні автоматизованих методів оцінювання якості рефератів, а також у вивченні можливостей адаптації технології до вузькоспеціалізованих наукових корпусів і довгих документів зі складною внутрішньою структурою.

References

- Barrios, F., López, F., Argerich, L., & Wachenchauser, R. (2016). Variations of the similarity function of TextRank for automated summarization. *arXiv preprint*. <https://doi.org/10.48550/arXiv.1602.03606>
- Beltagy, I., Peters, M., & Cohan, A. (2020). Longformer: The long-document transformer. *arXiv preprint*. <https://doi.org/10.48550/arXiv.2004.05150>
- Cao, Y., Dong, D., & Wei, F. (2024). Multilingual summarization with transformer models: A survey. *ACM Computing Surveys*, 56(10), article ID 248. <https://doi.org/10.1145/3673030>
- Dhaini, M., Erdogan, E., Bakshi, S., & Kasneci, G. (2024). *Explainability meets text summarization: A survey*. In *Proceedings of the 17th International Conference on Natural Language Generation* (pp. 631–645). <https://doi.org/10.18653/v1/2024.inlg-main.49>
- Dong, Y., Shen, Y., & Zhang, M. (2024). Scientific document summarization: A survey. *Information Fusion*, 104, article ID 102145. <https://doi.org/10.1016/j.inffus.2023.102145>
- Dyak, T. P., & Hrytsiuk, Y. I. (2024). Application of corpus tools to identify keywords of Ukrainian rebel songs as a genre of the folklore discourse. *Scientific Bulletin of UNFU*, 34(7), 60–71. <https://doi.org/10.36930/40340708>
- Erkan, G., & Radev, D. R. (2004). LexRank: Graph-based lexical centrality AS salience in text summarization. *Journal of Artificial Intelligence Research*, 22, 457–479. <https://doi.org/10.1613/jair.10396>
- Hrysai, S. R., & Dyak, T. P. (2025). Using digital corpus tools to study gendered aspects of political speeches. *Scientific Bulletin of UNFU*, 35(5), 108–115. <https://doi.org/10.36930/40350512>
- Koto, F., & Lau, J. H. Y. (2024). Large language models for meta-evaluation of summaries. In *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics* (pp. 1–15). <https://doi.org/10.18653/v1/2024.naacl-long.1>
- Kryściński, W., McCann, B., Xiong, C., & Socher, R. (2020). Evaluating the factual consistency of abstractive text summarization. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing* (pp. 9332–9346). <https://doi.org/10.18653/v1/2020.emnlp-main.750>
- Kunanets, N., & Yaromych, M. (2025). Extracting concepts from literary texts using large language models. *Bulletin of Science and Education*, 32(2), 343–357. [https://doi.org/10.52058/2786-6165-2025-2\(32\)-343-357](https://doi.org/10.52058/2786-6165-2025-2(32)-343-357)
- Laban, P., Hsu, C., Kottur, S., & Choi, Y. (2022). SummaC: Revisiting NLI-based models for inconsistency detection in summarization. *Transactions of the Association for Computational Linguistics*, 10, 163–177. <https://aclanthology.org/2022.tacl-1.10/>
- Lewis, M., Liu, Y., Goyal, N., Ghazvininejad, M., Mohamed, A., Levy, O., Stoyanov, V., & Zettlemoyer, L. (2020). BART: Denoising sequence-to-sequence pre-training for natural language generation, translation, and comprehension. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics* (pp. 7871–7880). <https://doi.org/10.18653/v1/2020.acl-main.703>
- Lewis, M., Perez, E., Liu, Y., Ghazvininejad, M., & Zettlemoyer, L. (2024). Pretrained multilingual models for cross-lingual summarization. *Transactions of the Association for Computational Linguistics*, 12, 1–20. https://doi.org/10.1162/tacl_a_00612
- Liu, J., He, M., & Zhang, L. (2023). An improved TextRank method based on K-means for document summarization. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.4399141>
- Liu, Y., Iter, D., Xu, Y., Wang, S., Xu, R., & Zhu, C. (2023). G-Eval: NLG evaluation using GPT-4 with better human alignment. *arXiv*. <https://doi.org/10.48550/arXiv.2303.16634>
- Maynez, J., Narayan, S., Bohnet, B., & McDonald, R. (2020). On faithfulness and factuality in abstractive summarization. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics* (pp. 1906–1919). <https://doi.org/10.18653/v1/2020.acl-main.173>
- Mihalcea, R., & Tarau, P. (2004). TextRank: Bringing order into texts. In *Proceedings of the 2004 Conference on Empirical Methods in Natural Language Processing* (pp. 404–411). <https://doi.org/10.3115/W04-3252>
- Nenkova, A., & McKeown, K. (2011). Automatic summarization. *Foundations and Trends in Information Retrieval*, 5(2–3), 103–233. <https://doi.org/10.1561/15000000015>
- Pasichnyk, V., & Yaromych, M. (2025). Automated generation of technical documentation in IT using large language models. *Studia Methodologica*, 59, 250–273. <https://doi.org/10.32782/2307-1222.2025-59-22>
- Pasichnyk, V., & Yaromych, M. (2025). Comparative analysis of large language models in solving linguistic tasks. *Current Issues of the Humanities. Linguistics. Literary Studies*, 89(1), 288–305. <https://doi.org/10.24919/2308-4863/89-1-41>
- Pasichnyk, V., & Yaromych, M. (2025). Features of genre classification of literature using large language models. *Folium*, 6, 132–143. <https://doi.org/10.32782/folium/2025.6.19>
- Pasichnyk, V., & Yaromych, M. (2025). Large language models and ontologies in philological research: An analytical review of sources. *Current Issues of the Humanities. Linguistics. Literary Studies*, 83(3), 236–250. <https://doi.org/10.24919/2308-4863/83-3-35>
- Pasichnyk, V., & Yaromych, M. (2025). Methods and tools of large language models for conceptualizing the subject area "artificial intelligence". In *Science in the Modern World: Innovations and Challenges: Proceedings of the IX International Scientific and Practical Conference* (pp. 353–358). URL: <https://sci-conf.com.ua/wp-content/uploads/2025/05/science-in-the-modern-world-innovations-and-challenges-15-17.05.2025.pdf>
- Pasichnyk, V., Yaromych, M., Pavliv, I., & Kunanets, N. (2025). Information technology for self-analysis of large language model parameters. *Scientific Bulletin of UNFU*, 35(5), 130–144. <https://doi.org/10.36930/40350515>
- Peters, U., & Chin-Yee, B. (2025). Generalization bias in LLM summaries of scientific text. *Royal Society Open Science*, 12(4), article ID 241776. <https://doi.org/10.1098/rsos.241776>
- Raffel, C., Shazeer, N., Roberts, A., Lee, K., Narang, S., Matena, M., Zhou, Y., Li, W., & Liu, P. J. (2020). Exploring the limits of transfer learning with a unified text-to-text transformer. *Journal of Machine Learning Research*, 21(140), 1–67. <https://doi.org/10.48550/arXiv.1910.10683>
- Scialom, T., Dray, P.-A., Lamprier, S., Piwowarski, B., Staiano, J., Wang, A., & Gallinari, P. (2021). QuestEval: Summarization evaluation using question generation and answering. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing* (pp. 6594–6604). <https://doi.org/10.18653/v1/2021.emnlp-main.529>
- See, A., Liu, P. J., & Manning, C. (2017). Get to the point: Summarization with pointer-generator networks. In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics* (pp. 1073–1083). <https://doi.org/10.18653/v1/P17-1099>
- Shen, Y., Li, M., et al. (2023). Hybrid extractive-abstractive summarization for scientific documents. *Advances in Neural Information Processing Systems*, 36, 12345–12360. <https://doi.org/10.1016/j.aej.2026.01.051>
- Wang, T., Sun, T., & Li, T. (2023). Graph-based methods for document summarization: A review. *Information Processing & Management*, 60(4), article ID 103312. <https://doi.org/10.1016/j.ipm.2023.103312>
- Zhang, J., & Zhao, W. (2025). Abstractive summarization with LLMs: Challenges and advances. *Neural Computation*, 37(1), 45–67. https://doi.org/10.1162/neco_a_00345
- Zhang, J., Zhao, Y., Saleh, M., & Liu, P. (2020). PEGASUS: Pre-training with extracted gap-sentences for abstractive summarization. In *Proceedings of the 37th International Conference on Machine Learning* (pp. 11328–11339). <https://doi.org/10.5555/3524938.3525989>
- Zhang, H., Yu, P., & Zhang, J. (2025). From statistical methods to large language models: A review of summarization paradigm shifts. *ACM Transactions on Intelligent Systems and Technology*, 16(2), article ID 43. <https://doi.org/10.1145/3731445>

INFORMATION TECHNOLOGY OF AUTOMATED REFERENCE OF SCIENTIFIC TEXTS USING LARGE LANGUAGE MODELS

A comprehensive information technology for automated summarization of scientific texts in the field of artificial intelligence is presented, integrating extractive methods for identifying content-bearing fragments, abstractive summarization models based on modern transformer architectures, and automated mechanisms for quality assessment of generated results. The technology is implemented as an integrated multilevel cycle of intelligent text processing, encompassing preliminary structural segmentation and document cleaning, the formation of compact semantic representations, multistage generative abstraction, and systematic verification of summaries using formal and semantic metrics. The proposed architecture ensures the coordinated integration of statistical, linguistic, and neural approaches, enabling the preservation of both factual accuracy and semantic coherence in generated summaries. Three experimental observations were conducted using the GPT-5, Grok-4, and Gemini 2.5 models on English-, German-, and Ukrainian-language scientific texts. The obtained results demonstrated a high level of preservation of the logical structure of scientific works, interdisciplinary conceptual links, and terminological accuracy, and also confirmed the multilingual capability and domain adaptivity of the developed information technology. GPT-5 ensured the greatest depth of conceptual abstraction and the integration of scientific ideas into coherent cognitive structures; Grok-4 demonstrated sensitivity to the philosophical – analytical style of German scientific texts and complex theoretical argumentation; and Gemini 2.5 exhibited high accuracy and stylistic consistency when processing Ukrainian technical language and specialized scientific terminology. The proposed information technology enables the generation of coherent, informative, and factually accurate summaries across different languages and scientific disciplines, supports adaptation to diverse subject domains, and provides a solid foundation for its practical application in intelligent knowledge analysis systems, scientific repositories, educational platforms, digital libraries, and automated scientific literature review services aimed at supporting research activity, interdisciplinary integration, and the acceleration of scientific communication, as well as the development of next-generation cognitive and analytical infrastructures for academic knowledge processing.

Keywords: automated summarization; large language models; extractive text reduction; abstract text generalisation; multilingual scientific texts.