

**Рувінська В.М.**

Національний університет «Одеська політехніка»

**Максिमичев А.М.**

Національний університет «Одеська політехніка»

## КОЛАБОРАТИВНО-КОНТЕНТНА ФІЛЬТРАЦІЯ ДЛЯ РЕКОМЕНДАЦІЙНИХ СИСТЕМ З МЕТОЮ МОДЕЛЮВАННЯ КОНТЕКСТУ НА БАЗІ ГЛИБИННОГО НАВЧАННЯ

Досліджуються сучасні підходи до побудови рекомендацій, інтеграція підходів колаборативної та контентної фільтрації шляхом застосування моделей глибокого навчання для контекстно-залежної персоналізації. Проведено комплексний аналіз традиційних методів рекомендацій, таких як колаборативна фільтрація та контентна фільтрація, охоплюючи їхні етапи разом із притаманними їм проблемами, такими як розрідженість даних, проблеми холодного старту та упередженість популярності. Зокрема, увага приділяється ключовим показникам подібності, а також методам прогнозування, включаючи середньозважене значення та агрегацію подібних елементів. Обговорюються обмеження автономних моделей контентної та колаборативної фільтрації, такі як надмірна спеціалізація та недостатня продуктивність у динамічних або розріджених середовищах. Представлено новий підхід *Deep Feature-Enriched Context-Aware Hybrid Filtering (DFECA-HF)*, що інтегрує інформацію про взаємодії користувача з елементами, атрибути контенту та контекстну інформацію, таку як час, пристрій, місцезнаходження. Архітектура системи складається з трьох основних модулів: двошарової системи вбудовування (колаборативна та контентна фільтрація), модуля адаптації контексту, який динамічно налаштовує профілі користувачів за допомогою послідовних моделей, таких як GRU та TCN, та адаптивного механізму зважування, який балансує виходи контентного, колаборативного та контекстного модулів на основі повноти даних та поведінки користувача. Майбутні дослідження мають бути зосереджені на інтеграції навчання з підкріпленням для довгострокової взаємодії з користувачами, застосуванні графових нейронних мереж для фіксації складних взаємозв'язків між користувачем та елементом, а також включенні мультимодальних даних, таких як зображення, аудіо, текст, для подальшого збагачення рекомендацій.

**Ключові слова:** рекомендаційні системи, колаборативна фільтрація, контентна фільтрація, глибоке навчання, контекстно-залежна персоналізація.

**Постановка проблеми.** У сучасному світі цифрових технологій, рекомендаційні системи відіграють вирішальну роль у персоналізації користувацького досвіду, пропонуючи елементи, які найбільш відповідають індивідуальним інтересам та вподобанням. Від онлайн-магазинів та стрімінгових сервісів до соціальних мереж та інформаційних платформ, здатність ефективно рекомендувати релевантний зміст користувачам є ключовою для залучення та утримання аудиторії. Використання передових методів машинного навчання, як-от колаборативна фільтрація, контентна фільтрація, дозволяє системам рекомендацій дедалі точніше прогнозувати потреби та вподобання користувачів. Незважаючи на значний прогрес у цій області, існують численні

виклики та проблеми, які потребують подальших досліджень та інноваційних підходів – проблеми холодного старту, розрідженості даних, масштабованості [1, с. 1-35]. Ця стаття покликана дослідити основні типи систем рекомендацій, їхні особливості, переваги, та обмеження, визначити поточні тенденції, перспективи розвитку в даній галузі, а також запропонувати нові підходи.

**Аналіз останніх досліджень і публікацій.** Розглянуто задачі колаборативної фільтрації: збір даних, створення матриць, розрахунок подібності, генерація прогнозів, обробка розрідженості та масштабованості та проблеми колаборативної фільтрації. За темою збору даних розглянуто джерела збору відгуків та підходи для їх опрацювання [2]. За темою створення матриць досліджено

методи побудови матриць взаємодій [3]. Для теми розрахунку подібності розглянуто методи розрахунку подібності Косинусна подібність, коефіцієнт кореляції Пірсона та індекс Жаккара [4-6]. За темою генерації прогнозів досліджено методи побудови прогнозів Середньозважене значення та Агрегування подібних елементів [7]. Для теми обробки розрідженості та масштабованості розглянуто метод подолання проблеми розрідженості Singular Value Decomposition (SVD) [8]. За темою проблем колаборативної фільтрації розглянуто проблеми розрідженості даних, Scalability problem, Cold Start Problem, First Rater problem, Gray Sheep problem, Popularity Bias та Synonymy problem [9-12].

Розглянуто задачі та складові контентної фільтрації: профілі елементів, профілі користувачів та проблеми контентної фільтрації. За темою профілів елементів досліджено характеристики елементів та їх вплив на архітектуру системи. Для теми профілів користувачів розглянуто профілі користувачів, методи їх побудови [13]. За темою проблем контентної фільтрації розглянуто проблеми обмеження атрибутів елементів, *надмірної спеціалізації* та холодного старту для нових елементів [14].

**Метою дослідження** є вивчення різних методів для створення рекомендацій, зокрема колаборативної фільтрації та контентної фільтрації, для визначення їхньої ефективності, точності та придатності в різних застосуваннях. Дослідження має на меті порівняти ці методи з метою виявлення їх сильних та слабких сторін, а також запропонувати засоби для вирішення проблем.

#### **Виклад основного матеріалу дослідження.**

1. Колаборативна фільтрація. Збір даних є першим важливим етапом у колаборативній фільтрації, оскільки якість та кількість зібраних даних безпосередньо впливають на ефективність рекомендаційної системи. Цей процес включає збір зворотного зв'язку у вигляді явних та неявних відгуків. Явними відгуками є прямі вказівки користувачів щодо їхніх уподобань, як правило у формі оцінок (як то рейтинг в зірках для фільмів чи продуктів) або бінарних оцінок "подобається" чи "не подобається". Явні відгуки є дуже інформативними, але їх отримання може бути ускладненим, оскільки формування явних відгуків вимагає активної участі користувача. Для збору явних відгуків зазвичай використовують опитування та анкети, які дають детальну інформацію про вподобання та задоволеність користувачів. Альтернативним підходом є збір неявних відгуків, які включають непрямі вказівки на вподобання

користувача, такі як історія перегляду, записи про покупки, час перегляду та рейтинг кліків. Неявні відгуки є більш численними, ніж явні, але менш інформативними. Збір неявних відгуків проводиться під час використання веб-сайтів і програм за допомогою файлів cookie, засобів відстеження сеансів і файлів журналів. Також джерелом неявних відгуків є журнали транзакцій – записи про взаємодію з платформами електронної комерції дають змогу зрозуміти вподобання користувачів [2, с. 173-197].

1.1 Побудова матриці взаємодій. Після збору даних їх потрібно структурувати у формі, яку алгоритми можуть ефективно обробляти, це досягається шляхом побудови матриці елементів користувача. Кожен рядок представляє користувача, кожен стовпець представляє елемент, а записи в матриці представляють взаємодію між користувачами та елементами. Цією взаємодією можуть бути оцінки, покупки, перегляди або інша форма зворотного зв'язку. Якщо дані містять явні оцінки, ці оцінки безпосередньо розміщуються в матриці. Також оцінки можуть бути представлені як двійкові дані, де 1 означає взаємодію, а 0 означає відсутність взаємодії. У деяких випадках для заповнення матриці можна використовувати частоту взаємодій, надаючи більшу вагу елементам, з якими користувачі взаємодіяли частіше. Різним типам взаємодій можна присвоїти різні ваги залежно від їх ступеня важливості. Так, замовлення може мати більшу задану вагу, ніж перегляд або клік. Якщо деякі користувачі схильні надавати вищі загальні оцінки, їх дані можуть бути нормалізовані за допомогою методу центрування середнього значення, коли середня оцінка кожного користувача віднімається від їхніх оцінок. Уподобання користувачів і популярність предметів можуть змінюватися з часом. В цьому випадку доцільним є включення часової динаміки в матрицю взаємодій. Це може передбачати додавання функцій затухання з часом до старіших взаємодій або сегментування даних за часовими рамками. Для систем заснованих на неявних відгуках додаткову інформацію, таку як демографічні дані користувачів або метадані товару, можна інтегрувати в матрицю, щоб збагатити дані [3, с. 305].

1.2 Розрахунок подібності є компонентом колаборативної фільтрації, що дозволяє системі ідентифікувати та кількісно оцінювати зв'язки між користувачами або елементами на основі їх моделей взаємодії. Цей розрахунок керує процесом рекомендацій, визначаючи, наскільки різні користувачі або елементи схожі один на одного.

Колаборативна фільтрація на основі користувачів є підходом, який обчислює схожість між користувачами на основі їхніх профілів взаємодії. Ідея полягає в тому, щоб ідентифікувати користувачів, які мають подібні вподобання, і використовувати їхні рейтинги для прогнозування оцінок для елементів, з якими цільовий користувач ще не взаємодіяв. Колаборативна фільтрація на основі елементів є підходом, у якому подібність між елементами обчислюється на основі того, як користувачі з ними взаємодіяли. Так, два предмети, які часто оцінюють ті самі користувачі, вважаються схожими. Рекомендації створюються на основі елементів, схожих на ті, які користувачеві сподобалися в минулому [4]. Для визначення рівня подібності використовуються різні методи в залежності від типу оцінок.

*Косинусна подібність* – це метод порівняння користувачів або елементів шляхом вимірювання косинуса кута між їхніми векторами ознак у багатовимірному просторі. У фільтрації на основі користувачів вектор користувача включає оцінки для всіх елементів, тоді як у фільтрації на основі елементів вектор елемента включає всі оцінки користувачів. Подібність обчислюється як скалярний добуток двох векторів, поділений на добуток їхніх величин (формула 1).

$$sim_{cos} = \frac{A \cdot B}{\|A\| \|B\|} \quad (1)$$

де  $A$ ,  $B$  – вектори ознак елементів,  $\|A\|$ ,  $\|B\|$  – довжини векторів. Результат коливається від -1 до 1, але в колаборативній фільтрації лише діапазон від 0 до 1 зазвичай має значення через невід’ємні дані. Косинусна подібність особливо ефективна для порівняння розріджених векторів і зосереджена на орієнтації вектора, а не на величині (кількість елементів, оцінених користувачем є менш актуальною, ніж самі оцінені елементи), що робить її добре придатною для завдань рекомендацій [5, с. 1].

*Коефіцієнт кореляції Пірсона* вимірює лінійну залежність між двома змінними, яку зазвичай використовують у колаборативній фільтрації для порівняння моделей оцінювання користувачів або елементів. Коефіцієнт враховує індивідуальні тенденції оцінювання, віднімаючи середню оцінку від кожного значення, зосереджуючись на відхиленнях від середньої поведінки (формула 2):

$$sim_p = \frac{\sum_{i=0}^n (A_i - \bar{A})(B_i - \bar{B})}{\sqrt{\sum_{i=0}^n (A_i - \bar{A})^2} \sqrt{\sum_{i=0}^n (B_i - \bar{B})^2}}, \quad (2)$$

де  $A_i$ ,  $B_i$  – компоненти векторів  $A$ ,  $B$ , а  $\bar{A}$  та  $\bar{B}$  – середні значення векторів відповідно. Результат

коливається від -1 до 1, що вказує на негативну (-1), відсутність (0) або позитивну лінійну кореляцію (1). У фільтрації на основі користувачів коефіцієнт це допомагає ідентифікувати користувачів зі схожою поведінкою оцінювання, тоді як у фільтрації на основі елементів це ідентифікує елементи, оцінені користувачами однаково.

*Індекс Жаккара* вимірює подібність між множинами та корисний у колаборативній фільтрації бінарних даних. Він розраховується як відношення розміру перетину двох наборів до розміру їх об’єднання (формула 3):

$$sim_j = \frac{|A \cap B|}{|A \cup B|}. \quad (3)$$

Його значення коливається від 0 (немає збігів) до 1 (набори ідентичні). Індекс ефективний для розріджених даних, оскільки враховує лише наявність взаємодій, а не їх частоту чи величину. Він обчислювально простий і допомагає ідентифікувати користувачів або елементи зі схожими моделями взаємодії, що робить його придатним для великомасштабних рекомендаційних систем, де тип взаємодії має більше значення, ніж інтенсивність [6, с. 3].

**1.3 Генерація прогнозів.** На цьому етапі оцінки подібності використовуються для прогнозування того, як користувачі можуть оцінювати елементи, з якими вони ще не стикалися, використовуючи такі методи, як середньозважене значення та агрегування подібних елементів.

*Середньозважене значення.* Цей метод оцінює рейтинг користувача, враховуючи середньозважене значення оцінок подібних користувачів (фільтрація на основі користувачів) або подібних елементів, оцінених користувачем (фільтрація на основі елементів).

*Агрегування подібних елементів.* Цей метод, поширений у фільтрації на основі елементів, прогнозує рейтинг користувача шляхом агрегування його оцінок для подібних елементів, зважених за подібністю. Він припускає, що користувачі оцінюють подібні елементи подібним чином, ефективно використовуючи минулу поведінку для прогнозування уподобань [7, с. 780].

**1.4 Обробка розрідженості та масштабованість** у системах колаборативної фільтрації має важливе значення для підтримки продуктивності, оскільки розмір набору даних збільшується, а матриця взаємодії між користувачем і елементом залишається в основному незаповненою. Існують наступні стратегії вирішення зазначених проблем: матрична факторизація, методи застосо-



вувати до сусідства, введення порогових значень розрідженості, кластеризація.

1.5 Матрична факторизація – це набір технік які вирішують як проблеми розрідженості, так і розмірності, властиві великим матрицям взаємодії між користувачем і елементом. Завдяки розкладанню великої розрідженої матриці на більш щільні матриці нижчої розмірності факторизація сприяє більш керованому та глибокому представленню базових даних. Основним методом матричної факторизації є SVD.

*Декомпозиція сингулярного значення (SVD)* – це фундаментальний метод факторизації матриці, який широко використовується в галузі колаборативної фільтрації та науки про дані. SVD розкладає матрицю  $A$  на три матриці:  $U$ ,  $\Sigma$  і  $V$ . Математично це виражається так (формула 4):

$$A = U \Sigma V^T, \quad (4)$$

де  $A$  – початкова матриця оцінки елементів користувача,  $U$  (ліві сингулярні вектори) представляє матрицю, пов'язану з користувачем,  $\Sigma$  (сингулярні значення) – діагональна матриця сингулярних значень, упорядкованих у порядку спадання, а  $V^T$  (праві сингулярні вектори) – матриця, пов'язана з елементами.  $U$  містить ознаки користувача, де кожен стовпець представляє приховану ознаку.  $\Sigma$  кількісно визначає силу кожної прихованої ознаки.  $V^T$  містить характеристики елементів, що відповідають стовпцям у  $U$ .

У колаборативній фільтрації не всі характеристики однаково корисні – SVD дозволяє зменшити розмірність шляхом вибору топ  $k$  сингулярних значень та їх відповідних векторів у  $U$  та  $V^T$ , ефективно захоплюючи найважливіші базові шаблони, одночасно відкидаючи шум та менш важливі деталі. Матриці зменшеної розмірності використовуються для реконструкції наближення оригінальної матриці з нижчою розмірністю (формула 5):

$$A_k = U_k \Sigma_k V_k^T, \quad (5)$$

де  $A_k$  включає прогнози для відсутніх записів [8, с. 635].

1.6 Колаборативна фільтрація стикається з кількома властивими проблемами, які можуть вплинути на ефективність систем рекомендацій.

*Розрідженість даних.* Матриця взаємодії між користувачем і елементом у колаборативній фільтрації зазвичай розріджена, оскільки не всі користувачі оцінюють усі елементи. Ця розрідженість ускладнює ефективний пошук подібності між користувачами або елементами, оскільки може бути недостатньо збігів між елементами, з якими

взаємодіяли різні користувачі, або користувачами, які взаємодіяли з різними елементами.

*Scalability problem.* Зі збільшенням кількості користувачів і елементів обчислювальна складність алгоритмів колаборативної фільтрації може значно зрости. Масштабованість стає проблемою, особливо з методами на основі пам'яті, які вимагають попарних обчислень між усіма користувачами або всіма елементами [9, с. 608].

*Cold Start Problem.* Проблема холодного старту користувача виникає, коли нові користувачі приєднуються до платформи, тоді недостатньо даних про їхні вподобання, що ускладнює точну рекомендацію елементів, оскільки системі бракує історичних даних про взаємодію. Проблема холодного старту елемента виникає, коли нові елементи додаються до каталогу, тоді вони не мають оцінок або мають низьку оцінку, тому системі важко рекомендувати ці товари користувачам, оскільки існує небагато точок даних взаємодії для аналізу [10, с. 774].

*First Rater problem.* Проблема першого оцінювача у колаборативній фільтрації виникає, коли елемент був оцінений дуже малою кількістю користувачів, що призводить до недостатніх даних для точного обчислення подібності або створення надійних прогнозів. Через брак даних важко рекомендувати продукт іншим користувачам, підвищується вразливість до шуму та відхилень від обмежених доступних оцінок.

*Gray Sheep problem.* Користувачі, чий вподобання не збігаються з певною групою користувачів або тенденцією (сірі вівці), створюють проблему для систем колаборативної фільтрації, які покладаються на пошук подібних користувачів або елементів. Ці користувачі можуть виявити, що рекомендації системи є менш актуальними [11].

*Popularity Bias.* Системи колаборативної фільтрації, як правило, рекомендують елементи, які вже популярні та часто оцінюються. Це упередження може призвести до відсутності різноманітності в рекомендаціях і може перешкоджати наданню рекомендацій менш популярних або вузько спеціалізованих елементів, навіть якщо вони можуть зацікавити певних користувачів.

*Synonymy problem.* Різні предмети, які по суті схожі, можуть не розпізнаватися системою як такі, якщо вони часто не оцінюються одними і тими самими користувачами. Це може призвести до надлишкових рекомендацій і нездатності консолідувати налаштування користувача для подібних елементів [12].

2. Контентна фільтрація – це метод, який використовується в системах рекомендацій, щоб про-

понувати користувачам елементи на основі атрибутів елементів і профілю уподобань користувача. Цей метод передбачає збір характеристик для профілей користувача та елементів.

2.1 Профілі елементів є важливими для забезпечення релевантних рекомендацій, оскільки вони визначають ключові характеристики кожного елемента. Створення цих профілів починається з вибору відповідних атрибутів, які відрізняються залежно від предметної області, наприклад, жанру, режисера та рейтингів для рекомендацій фільмів. Дані для цих атрибутів збираються через API, бази даних або веб-краулери, а іноді й за допомогою ручних анотацій для складних ознак, таких як настрої чи стиль. Після зібраних даних обробляються – це може включати одноразове кодування категоріальних даних, нормалізацію числових значень або використання складніших методів, таких як вбудовування векторів, щоб зафіксувати нюанси більш складних атрибутів. Оскільки атрибути елементів можуть змінюватися з часом, профілі необхідно регулярно оновлювати. Контентна фільтрація порівнює профілі користувачів та елементів, використовуючи такі міри подібності, як косинусна подібність або евклідова відстань, генеруючи прогнози, що використовуються для ранжування та рекомендації найбільш релевантних елементів. [13].

2.2 Профілі користувачів мають вирішальне значення для надання персоналізованих рекомендацій. Їх створення починаються зі збору даних, дані можуть бути як явними (оцінки та вподобання користувачів), так і неявними (кліки чи історія переглядів). На основі цього ключові атрибути, такі як улюблені жанри чи бренди, визначаються та кількісно оцінюються за допомогою таких показників, як середня оцінка чи частота взаємодій. Потім дані перетворюються у формат векторів вподобань. Оскільки інтереси користувачів змінюються з часом, профілі необхідно регулярно оновлювати, включаючи нові взаємодії та відкидаючи застарілі дані. Рекомендації генеруються шляхом порівняння профілів користувачів та елементів за допомогою показників подібності, таких як косинусна подібність або кореляція Пірсона, що забезпечує актуальність результатів з часом [13].

2.3 Контентна фільтрація стикається з кількома проблемами, які впливають на якість рекомендацій.

*Обмеження атрибутів елементів* виникає тому, що системи покладаються лише на попередньо визначені атрибути, які можуть не повністю враховувати вподобання користувача. Фактори,

які важко оцінити кількісно, такі як теми чи стиль, можуть бути пропущені, а відсутні або неточні дані можуть призвести до поганих рекомендацій.

*Надмірна спеціалізація* виникає, коли системи неодноразово пропонують подібні елементи, обмежуючи різноманітність і створюючи бульбашки фільтрів, коли користувачеві рекомендуються тільки елементи з його обмеженої сфери інтересів, що з часом може знизити залученість.

*Холодний старт для нових елементів* стосується труднощів у рекомендації нових елементів через брак даних про взаємодію, що затримує їх включення до рекомендацій, навіть якщо вони є дуже релевантними [14].

3. Deep Feature-Enriched Context-Aware Hybrid Filtering (DFECA-HF). Пропонується новий підхід до рекомендацій – контекстна гібридна фільтрація, збагачена ознаками – який інтегрує колаборативну фільтрацію та контентну фільтрацію в рамках єдиної контекстно-залежної архітектури, дані користувачів та елементів збагачуються шляхом інтеграції кількох джерел, таких як поведінкові взаємодії, атрибути контенту та контекстна інформація, у інформативніші представлення за допомогою моделей глибинного навчання. Підхід запропонований для усунення ключових обмежень традиційних систем рекомендацій, таких як розрідженість даних, холодний старт, popularity bias та надмірна спеціалізація, одночасно покращуючи персоналізацію та релевантність завдяки динамічному включенню контекстної інформації.

3.1 Архітектура DFECA-HF. Пропонується модель, що складається із трьох основних модулів: двошарової системи вбудовування, модуля контекстної адаптації та механізму адаптивного зважування. Ці компоненти працюють разом для вилучення, збагачення та поєднання інформації користувача, елемента та контексту для персоналізованих рекомендацій.

Двошарова система вбудовування поєднує колаборативні та контентні представлення.

*Колаборативний рівень* – використовує методи матричної факторизації (розкладання сингулярних значень, SVD) або моделі глибинного навчання (нейронна колаборативна фільтрація) для визначення прихованих представлень з матриці взаємодії користувач-елемент. Ці вбудовування фіксують неявні уподобання користувачів, отримані з історичних взаємодій.

*Контентний рівень* – отримує семантичні представлення зі структурованих та неструктурованих атрибутів елементів та користувачів (жанрів, описів, демографічних даних) за допомогою

глибинних кодерів, таких як згорткові нейронні мережі (CNN), рекурентні нейронні мережі (RNN) або моделі на основі трансформерів. Ці представлення фіксують явні ознаки, які особливо корисні в сценаріях холодного запуску.

Вихідні дані з обох шарів об'єднуються за допомогою механізму нелінійного комбінування (багатошаровий перцептрон або об'єднання на основі уваги), що призводить до єдиного представлення, яке фіксує як поведінкові моделі, так і семантику на рівні атрибутів.

**Механізм контекстної адаптації.** Для підвищення часової та ситуативної релевантності рекомендацій модель включає модуль контекстної адаптації. Цей компонент моделює динамічну поведінку користувача з часом та в різних сценаріях використання (тип пристрою, час доби, місцезнаходження). Нещодавні послідовності взаємодії кодується за допомогою послідовних моделей, таких як ґратовані рекурентні одиниці (GRU) або часові згорткові мережі (TCN), які дозволяють системі оновлювати профілі користувачів у режимі реального часу. Контекстуальні вбудовування потім інтегруються в кінцеве представлення користувача перед обчисленням подібності.

**Механізм адаптивного зважування.** Кінцевий бал рекомендації обчислюється за допомогою функції адаптивного зважування, яка динамічно коригує внесок кожного модуля – контентного (CF), колаборативного (CBF) та контекстного (CTX) (формула 6):

$$Score_{ui} = \alpha_u \cdot CF_{ui} + \beta_u \cdot CBF_{ui} + \gamma_u CTX_{ui}, \quad (6)$$

де  $\alpha_u, \beta_u, \gamma_u \in [0, 1]$  – це ваги, залежні від користувача, які визначаються під час навчання моделі. Ці ваги модулюються на основі повноти даних, тривалості історії взаємодії, залученості користувача та ситуативного контексту. Так, для нового користувача без історії взаємодії модель надаватиме пріоритет інформації контентного та контекстного модулів.

3.2 Запропонований підхід DFECa-HF має наступні переваги.

**Стійкість до холодного старту** – поєднує особливості контенту та контексту для створення рекомендацій за відсутності історичних даних.

**Контекстно-залежна персоналізація** – динамічно адаптується до ситуативної інформації, такої як час та платформа, підвищуючи задоволеність користувачів.

**Покращена різноманітність** – зменшує надмірну спеціалізацію шляхом балансування кількох джерел даних та впровадження нових елементів.

**Масштабованість** – архітектура на основі вбудовування дозволяє проводити офлайн-навчання та ефективну оцінку результатів, що робить її придатною для великомасштабних систем.

**Висновки.** В ході цього дослідження було проведено аналіз різноманітних методів рекомендацій, зокрема колаборативної фільтрації, контентної фільтрації. Виявлено ключові переваги та недоліки кожного методу, а також їх придатність до різних контекстів застосування.

Колаборативна фільтрація є ефективною для систем, які містять велику кількість даних про взаємодії користувачів. Цей метод страждає від проблеми холодного старту та розрідженості даних.

Контентна фільтрація надає стабільність у рекомендаціях, використовуючи ознаки елементів, але може обмежувати відкриття нових інтересів користувачем через схожість рекомендованого контенту.

У статті пропонується гібридний підхід – Deep Feature-Enriched Context-Aware Hybrid Filtering (DFECa-HF), який поєднує переваги методів колаборативної та контентної фільтрації та доповнює їх контекстною інформацією за допомогою моделей глибинного навчання. Запропонована архітектура дозволяє адаптувати рекомендації до поведінки користувачів та умов взаємодії, що сприяє покращенню персоналізації, релевантності та різноманітності результатів. Показано можливі варіанти використання цього підходу в застосунках для рекомендації фільмів.

Подальші дослідження будуть зосереджені на інтеграції графових нейронних мереж, мультимодальних даних та методів навчання з підкріпленням для довгострокової оптимізації рекомендацій.

### Список літератури:

1. Francesco Ricci, Lior Rokach, Bracha Shapira. Recommender systems handbook: second edition. Springer US. 2015. pp. 1003.
2. Walid Maalej, Thomas Fritz, and Romain Robbes. Collecting and processing interaction data for recommendation systems. *Recommendation Systems in Software Engineering*. 2014. P. 173–197
3. Rui-sheng Zhang, Qi-dong Liu, Chun-Gui, Jia-Xuan Wei, Huiyi-Ma Collaborative Filtering for Recommender Systems. Second International Conference on Advanced Cloud and Big Data, Huangshan, China. 2014. P. 301–308
4. Halime Khojamli, Jafar Razmara. Survey of similarity functions on neighborhood-based collaborative filtering. *Expert Systems with Applications*. 2021. vol. 185, DOI: <https://doi.org/10.1016/j.eswa.2021.115482>



5. Najim Dehak, Reda Dehak, James Glass, Douglas Reynolds, and Patrick Kenny. Cosine Similarity Scoring without Score Normalization Techniques. URL: [https://groups.csail.mit.edu/sls/publications/2010/Dehak\\_Odyssey.pdf](https://groups.csail.mit.edu/sls/publications/2010/Dehak_Odyssey.pdf)
6. Fletcher S., Islam M. Z. Comparing sets of patterns with the Jaccard index. *AJIS*. 2018. vol. 22. DOI: <https://doi.org/10.3127/ajis.v22i0.1538>
7. Beliakov G., Calvo T. & James S. Aggregation functions for recommender systems. *Recommender Systems Handbook, Second Edition*. 2011. P. 777–808. DOI: [https://doi.org/10.1007/978-1-4899-7637-6\\_23](https://doi.org/10.1007/978-1-4899-7637-6_23)
8. Barathy R., Chitra P. Applying Matrix Factorization In Collaborative Filtering Recommender Systems / 2020 6th International Conference on Advanced Computing and Communication Systems (ICACCS), Coimbatore, India. 2020. P. 635–639. DOI: <https://doi.org/10.1109/ICACCS48705.2020.9074227>
9. Xinchang K., Park D. S., Vilakone P. Movie Recommendation System Using Social Network Analysis and k-Nearest Neighbor. *Lecture Notes in Electrical Engineering*. 2020. P. 606–611. DOI: [https://doi.org/10.1007/978-981-13-9341-9\\_104](https://doi.org/10.1007/978-981-13-9341-9_104)
10. Kim H. N., Ha I., Lee K. S., Jo G. S., El-Saddik A. Collaborative user modeling for enhanced content filtering in recommender systems / *Decision Support Systems*. 2011. vol. 51. P. 772–781. DOI: <https://doi.org/10.1016/j.dss.2011.01.012>
11. Barragáns-Martínez A. B., Costa-Montenegro E., Burguillo J. C., Rey-López M., Mikic-Fonte F. A., Peleteiro A. A hybrid content-based and item-based collaborative filtering approach to recommend TV programs enhanced with singular value decomposition. *Information Sciences*. 2010. vol. 180. P. 4290–4311. DOI: <https://doi.org/10.1016/j.ins.2010.07.024>
12. Patel B., Desai P. & Panchal U. Methods of recommender system: A review. *Proceedings of 2017 International Conference on Innovations in Information. Embedded and Communication Systems*. 2017. P. 1–4. DOI: <https://doi.org/10.1109/ICIIECS.2017.8275856>
13. Sánchez P., Bellogín A. Building user profiles based on sequences for content and collaborative filtering. *Information Processing & Management*. 2019. vol. 56. P. 192–211. DOI: <https://doi.org/10.1016/j.ipm.2018.10.003>
14. Shah K., Salunke A., Dongare S., Antala K. Recommender systems: An overview of different approaches to recommendations. *Proceedings of International Conference on Innovations in Information. Embedded and Communication Systems*. 2017. DOI: <https://doi.org/10.1109/ICIIECS.2017.8276172>

#### **Ruvinska V.M., Maksimych A.M. COLLABORATIVE CONTENT FILTERING FOR RECOMMENDER SYSTEMS FOR CONTEXT MODELING BASED ON DEEP LEARNING**

*The paper explores advanced approaches to building recommendations, integrating collaborative and content filtering approaches by applying deep learning models for context-aware personalization. The comprehensive analysis is conducted which includes traditional recommendation methods such as collaborative filtering and content filtering, covering their stages – data collection, matrix creation, similarity computation and prediction generation – along with their inherent problems such as data sparsity, cold start problems and popularity bias. In particular, attention is paid to key similarity metrics such as cosine similarity, Pearson correlation coefficient and Jaccard index, as well as prediction methods including weighted average and aggregation of similar items. The limitations of autonomous content and collaborative filtering models, such as over-specialization and insufficient performance in dynamic or sparse environments, are discussed. Then a new approach is presented, Deep Feature-Enriched Context-Aware Hybrid Filtering (DFECA-HF). This proposed approach integrates user interaction data with elements, content attributes, and contextual information such as time, device, and location. The architecture consists of three main modules: a two-layer embedding system (collaborative and content filtering), a context adaptation module that dynamically adjusts user profiles using sequential models such as GRU and TCN, and an adaptive weighting mechanism that balances the outputs of the content, collaborative, and context modules based on data completeness and user behavior. Future research will focus on integrating reinforcement learning for long-term user interactions, applying graph neural networks to capture complex user-item relationships, and incorporating multimodal data such as images, audio, text to further enrich recommendations.*

**Key words:** recommender systems, collaborative filtering, content filtering, deep learning, context-aware personalization.

Дата надходження статті: 24.09.2025

Дата прийняття статті: 13.10.2025

Опубліковано: 16.12.2025