

КОПИЛЬЧАК ОЛЕГ

Національний університет «Львівська політехніка»

<https://orcid.org/0009-0009-3295-6887>e-mail: oleh.a.kopylchak@lpnu.ua**КАЗИМИРА ІРИНА**

Національний університет «Львівська політехніка»

<https://orcid.org/0000-0003-1597-5647>e-mail: iryna.y.kazymyra@lpnu.ua

СУЧАСНІ ПІДХОДИ ДО ПЕРСОНАЛІЗОВАНОГО ПЛАНУВАННЯ НАВЧАЛЬНИХ ТРАЄКТОРІЙ: ГРАФИ ЗНАНЬ, КОГНІТИВНА ДІАГНОСТИКА, НЕЙРОННІ МЕРЕЖІ ТА ВЕЛИКІ МОВНІ МОДЕЛІ

У даному дослідженні здійснено систематизований критичний огляд сучасних підходів до персоналізованого планування навчальних траєкторій у цифровій освіті, де особлива увага приділена методам, що враховують індивідуальні відмінності студентів, їхній поточний рівень знань та освітні цілі. Потреба у персоналізації зумовлена зростанням масштабів онлайн-навчання, високою варіативністю рівня підготовки слухачів та різноманітністю форматів навчального контенту. Проаналізовано шість ключових груп підходів, які відображають еволюцію досліджень у цій сфері. Графові моделі знань, зокрема контекстуальні KG, забезпечують формальне представлення предметних областей, пререквізитів та семантичних зв'язків, створюючи основу для побудови когнітивно узгоджених навчальних шляхів. Методи когнітивної діагностики (DINA, V-DINA) дозволяють оцінювати рівень засвоєння конкретних понять і виявляти провали у знаннях, проте вимагають якісних і повних логів взаємодії. Комбінаторно-оптимізаційні схеми, такі як Fuzzy-CDF у поєднанні з алгоритмом Apriori та оптимізацією рою частинок (PSO), реалізують впорядкування навчальних понять і добір ресурсів, забезпечуючи баланс між пререквізитами, когнітивними можливостями й стилем навчання студента. Окремий пласт становлять нечітко-нейронні моделі (fuzzy-ANN), які ґрунтуються на стилях навчання за Колбом (Kolb's LSI) і поєднують нечіткі ваги з адаптивністю штучних нейронних мереж, що дає змогу враховувати індивідуальні освітні стратегії. Послідовні та багатозадачні архітектури (Seq2Seq, LSTM з non-repeat loss) орієнтовані на прогнозування подальших кроків навчання і водночас реалізують трасування знань, що дозволяє відстежувати ймовірність успішного засвоєння матеріалу. Перспективним напрямом виявилися мультимодальні підходи, які інтегрують різні типи даних (текст, відео, поведінкові сигнали), а також використання великих мовних моделей (LLM), що забезпечують семантичне збагачення графів знань, інтерпретацію природномовних запитів користувачів і пояснюваність рекомендацій.

У процесі аналізу підкреслено сильні та слабкі сторони кожного з підходів: від високої структурованості й пояснюваності графових методів до вимогливості нейронних мереж щодо обсягів даних та ресурсів; від точності когнітивної діагностики до її залежності від повноти логів; від гнучкості LLM до викликів, пов'язаних із ризиком галюцинацій та потребою в контролі якості. Огляд показав, що жоден із підходів не є універсальним і перспективним є створення гібридних архітектур, які поєднують структурованість графових моделей, діагностичні можливості когнітивних методів, адаптивність нейронних мереж і гнучкість LLM. Подальші дослідження варто зосередити на інтеграції мультимодальних сигналів, розробці протоколів відтворюваності експериментів, уніфікації метрик оцінювання та побудові масштабованих систем, здатних ефективно працювати в умовах великих освітніх середовищ (MOOCs).

Ключові слова: персоналізоване планування навчання, глибоке навчання, колаборативна фільтрація, рекомендаційна система.

KOPYLCHAK OLEH**KAZYMYRA IRYNA**

Lviv Polytechnic National University

MODERN APPROACHES TO PERSONALIZED LEARNING PATH PLANNING: KNOWLEDGE GRAPHS, COGNITIVE DIAGNOSIS, NEURAL NETWORKS, AND LARGE LANGUAGE MODELS

This study presents a systematic critical review of contemporary approaches to personalized learning path planning in digital education, with particular attention to methods that account for individual differences among students, their current knowledge levels, and their educational goals. The need for personalization is driven by the growing scale of online learning, the high variability in learners' preparedness, and the diversity of educational content formats. Six key groups of approaches are analyzed, reflecting the evolution of research in this domain. Knowledge graph models, particularly context-aware KGs, provide a formal representation of subject areas, prerequisites, and semantic relationships, thereby forming the basis for constructing cognitively coherent learning paths. Cognitive diagnostic methods (DINA, V-DINA) enable the assessment of mastery of specific concepts and the identification of knowledge gaps, although they require high-quality and comprehensive interaction logs. Combinatorial-optimization schemes, such as Fuzzy-CDF combined with the Apriori algorithm and particle swarm optimization (PSO), implement the sequencing of learning concepts and the selection of resources, ensuring a balance between prerequisites, cognitive capabilities, and the learner's preferred style. Another strand is represented by fuzzy-neural models (fuzzy-ANN), which build on Kolb's Learning Style Inventory (LSI) and integrate fuzzy weights with the adaptability of artificial neural networks, thus allowing for the consideration of individual educational strategies. Sequential and multitask architectures (Seq2Seq, LSTM with non-repeat loss) focus on predicting the learner's next steps while simultaneously implementing knowledge tracing, which makes it possible to monitor the probability of successful mastery of material. A promising direction is the use of multimodal approaches that integrate various types of data (text, video, behavioral signals), along with large language models (LLMs), which provide semantic enrichment of knowledge graphs, interpretation of natural language queries, and explainability of recommendations.

The analysis highlights the strengths and weaknesses of each approach: from the high degree of structure and explainability of graph-based methods to the data and resource demands of neural networks; from the accuracy of cognitive diagnostics to its dependence on log completeness; from the flexibility of LLMs to the challenges posed by hallucination risks and the need for quality control. The review shows that none of the approaches is universal, and the most promising direction lies in building hybrid architectures that combine the structured nature of knowledge graphs, the diagnostic capacity of cognitive methods, the adaptability of neural networks, and the flexibility of LLMs. Future research should focus

Постановка проблеми

Сучасна освіта характеризується стрімкими змінами та переходом до гнучких форм навчання, що потребує індивідуалізованого підходу до кожного учня або студента. Традиційні методи організації освітнього процесу, орієнтовані на середньостатистичного користувача, вже не можуть задовольнити всіх освітніх потреб і цілей окремого індивіда. Внаслідок цього зростає попит на персоналізацію освітніх траєкторій — підхід, який дозволяє адаптувати навчальні матеріали, їх послідовність і складність відповідно до особистих особливостей, запитів та навчальних цілей конкретного студента.

Сучасні підходи до персоналізації спираються на різні моделі представлення знань і студента: графи знань із пререквізитами та семантичними зв'язками; когнітивну діагностику (DINA/V-DINA) для оцінки опанування; аналітику взаємодій у реальному часі; послідовні моделі (Seq2Seq, LSTM) для прогнозування наступних кроків. Доповнюють їх системи підтримки добору та впорядкування контенту, що узгоджують матеріал із потребами користувача; на практиці застосовують колаборативну й контентну фільтрацію, неймережеві та гібридні підходи (див. [1], [2]). Ключові виклики — автоматичне виявлення зв'язків між поняттями, поєднання локальної й глобальної структури предметної області та робастність за обмежених або шумних даних. Перспективним є врахування мультимодальних сигналів (текст, відео, активність у середовищі) й узгодження плану з навчальними цілями, дедлайнами та часовими обмеженнями; додатково зростає потенціал великих мовних моделей (LLM, зокрема GPT) для інтерпретації запитів природною мовою і збагачення семантики навчальних об'єктів. Незважаючи на значні досягнення у розвитку технологій персоналізації, досі залишається актуальним питання ефективності застосування цих підходів у конкретних освітніх сценаріях, зокрема — у контексті планування індивідуальних освітніх траєкторій. Враховуючи широкий спектр існуючих методів та алгоритмів, виникає необхідність комплексного аналізу їх переваг, обмежень та можливостей інтеграції з великими мовними моделями.

Метою роботи - систематизований критичний огляд сучасних підходів до персоналізованого планування навчальних траєкторій, оцінка потенціалу неймережевих методів і LLM у моделюванні знань, когнітивній діагностиці та формуванні послідовностей, формування критеріїв відбору рішень і окреслення перспективних напрямів подальших досліджень.

Об'єктом дослідження є процеси та інформаційно-алгоритмічні засоби персоналізованого планування індивідуальних навчальних траєкторій у цифрових освітніх середовищах (MOOC/змішане навчання), зокрема на основі нейронних мереж і споріднених підходів: когнітивної діагностики, графових подань знань, послідовних і мультимодальних моделей, а також LLM.

Наукова новизна полягає в комплексному аналізі та оцінці можливостей інтеграції традиційних рекомендаційних алгоритмів і новітніх великих мовних моделей (LLM) саме в контексті персоналізованого освітнього планування. Зокрема, запропоновано аналіз перспективних шляхів поєднання цих технологій, що дозволяє поглибити наявні знання про ефективність застосування штучного інтелекту у сфері персоналізації навчання та окреслити нові перспективні напрямки подальших наукових досліджень у цій галузі.

Практична значущість отриманих результатів полягає у представленні комплексного аналізу ефективності існуючих алгоритмів і моделей, що допоможе освітнім установам та розробникам систем персоналізованого навчання обрати оптимальні підходи для реалізації персоналізованих навчальних середовищ. Результати аналізу можуть бути використані при проектуванні, розробці та впровадженні сучасних систем персоналізованого планування навчальних траєкторій, що підвищить якість та ефективність освітнього процесу.

Виклад основного матеріалу

Колаборативна фільтрація (Collaborative Filtering, CF) - один з найпоширеніших підходів, що базується на аналізі колективних уподобань користувачів. Ключова ідея полягає в тому, що користувачам, подібним за поведінкою, подобаються схожі елементи.

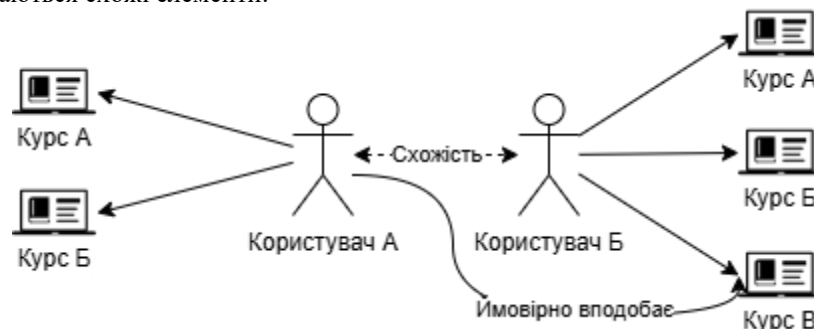


Рис. 1 Візуалізація принципу роботи колаборативної фільтрації: пошук схожих користувачів

Система знаходить «сусідів» цільового учня - інших студентів із подібними патернами вибору курсів і рекомендує йому курси, які ці схожі користувачі вже оцінили чи пройшли [3]. Для знаходження «сусідів» використовують функцію подібності. Часто для цієї задачі використовують коефіцієнт косинусної подібності:

$$\text{sim}(u, v) = \frac{\sum_{i \in I_{uv}} r_{ui} r_{vi}}{\sqrt{\sum_{i \in I_u} r_{ui}^2} \sqrt{\sum_{i \in I_v} r_{vi}^2}} \quad (1)$$

Де u, v - два користувачі, між якими обчислюється подібність; I_{uv} - множина спільно оцінених користувачами u, v курсів; r_{ui}, r_{vi} - оцінки, виставлені користувачами u, v для курсу i .

Наприклад, якщо студенти А і В пройшли багато однакових курсів, і В додатково прослухав курс Х, то Х може бути рекомендований студенту А. Item-item CF працює аналогічно, але шукає схожі курси: наприклад, якщо багато користувачів, подібних до вас, проходили курс Y одразу після курсу Z, то після Z система порадить Y.

Колаборативна фільтрація не вимагає аналізувати сам контент курсів - лише матрицю взаємодій «користувач-курс» (рейтинги, записи про завершення, перегляди тощо).

$$R = \begin{bmatrix} r_{11} & r_{12} & \dots & r_{1m} \\ r_{21} & r_{22} & \dots & r_{2m} \\ \dots & \dots & \dots & \dots \\ r_n & r_{n2} & \dots & r_{nm} \end{bmatrix} \quad (2)$$

Де R - матриця взаємодій користувачів (студентів) із курсами; r_{ij} - оцінка (рейтинг, перегляд, завершення) користувачем i курсу j .

Це дозволяє виявляти приховані зв'язки і несподівані рекомендації, коли курс популярний серед схожих студентів навіть без збігу тематики. У домені онлайн-освіти CF успішно застосовується для рекомендації MOOC-курсів за даними про реєстрації та успішність студентів[4]. Існують вдалі реалізації CF, наприклад, на основі матричного факторизації (Latent Factor models) та її нейромережових розширень. Так, метод Matrix Factorization будує приховані вектори для кожного учня і курсу, навчаючись відтворювати наявні рейтинги. Він ефективний і масштабується, але страждає від проблеми розрідженості даних - коли нових курсів і студентів дуже багато, а історія взаємодій надто мала для надійних прогнозів[3]. Сучасні підходи включають нейронні мережі для CF (наприклад, Neural CF, AutoRec), які замінюють лінійну модель факторизації багатопшаровою нейронною мережею, здатною моделювати нелінійні взаємозв'язки між користувачами і курсами. Це підвищує точність рекомендацій у разі складних залежностей, але водночас ускладнює інтерпретацію моделі.

Переваги CF у тому, що цей підхід враховує колективний досвід багатьох учнів, може рекомендувати популярні й високооцінені матеріали, не потребує ручної розмітки контенту. З недоліків можна виділити проблему нового користувача/курсу (cold start), відсутність рекомендацій для нетипових інтересів, можлива тенденція рекомендувати лише масово популярний контент. В освітньому контексті чиста CF не враховує навчальні цілі чи зміст - наприклад, вона може порадити курс, який «усім сподобався», але який не відповідає рівню знань конкретного студента. Тому на практиці CF часто поєднують з іншими методами в гібридні системи.

Контентно-орієнтовані рекомендації (Content-Based Filtering) ґрунтуються на аналізі властивостей самих навчальних матеріалів. Система намагається зрозуміти, який контент подобається користувачу, і підбирає схожий за змістом. Інакше кажучи, рекомендуються курси, подібні до тих, що вже переглядав чи високо оцінив студент [5].

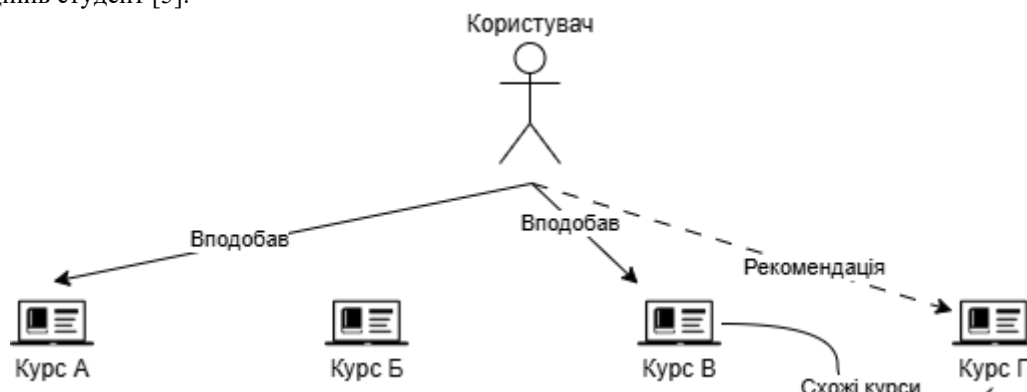


Рис. 2 Схема роботи контентно-орієнтованої рекомендаційної системи

В ролі контентних характеристик можуть виступати тема курсу, ключові слова з опису, рівень складності, тривалість, автор, формат (відео/текст) тощо. Для відеолекцій важливими ознаками є транскрипт або анотація відео, на основі яких можна визначити покриті тематики. Для текстових матеріалів (статей, підручників) застосовують NLP-аналіз: виділення ключових слів, категоризація по предметних областях, навіть визначення читацького рівня. Наприклад, якщо студент успішно пройшов курс «Вступ до програмування Python», контентний підхід може порекомендувати йому інші матеріали з Python або суміжні курси з

програмування, виходячи з схожих тегів і описів. Для визначення важливості слів у текстових матеріалах часто використовується TF-IDF (Term Frequency-Inverse Document Frequency):

$$TF - IDF(t, d) = TF(t, d) * \log \frac{N}{DF(t)} \quad (3)$$

Де, $TF(t, d)$ - частота появи терміна t в документі d ; N - загальна кількість документів (курсів чи лекцій); $DF(t)$ - кількість документів, у яких зустрічається термін t .

На відміну від колаборативної фільтрації, тут не потрібно багато користувачів - можна рекомендувати контент навіть першому студенту на платформі, якщо відомо, що його профіль чи початкові вподобання (введені теми цікавості) збігаються з атрибутами курсу. Це частково вирішує проблему cold start для нових елементів.

Сучасні контентні методи все більше використовують семантичні моделі глибокого навчання. Замість простого збігу ключових слів застосовуються попередньо навчені мовні моделі (наприклад, BERT) для перетворення опису курсу чи тексту лекції у векторні представлення (embeddings). Це дозволяє врахувати тонкі семантичні подібності: курс і стаття можуть не мати спільних слів, але модель зрозуміє, що обидва про алгоритми сортування. У статті [6] пропонується модель рекомендації курсів на основі BERT (Bidirectional Encoder Representations From Transformers), що інтегрує глибоке навчання і освітні дані.

Перевагами контентного підходу є прозорість (можна пояснити рекомендацію через спільні теми), відсутність залежності від чужих оцінок, придатність використання для нових курсів. З мінусів: обмежена різноманітність - система рекомендує дуже схожий контент, не пропонуючи нового, якість залежить від повноти і правильності метаданих (якщо опис курсу бідний, якість рекомендацій страждає), ігнорування суспільної думки (може радити матеріал, який відповідає темі, але нікому не сподобався). В освітніх платформах контентний підхід добре працює для каталогізації знань - наприклад, для створення списку всіх курсів з однієї тематики або знаходження відео, що покриває певне поняття, яке цікавить студента.

Гібридні системи поєднують кілька різних підходів, щоб взаємно компенсувати їх недоліки. У практичних реалізаціях дуже поширені саме гібридні стратегії, оскільки освітні потреби різноманітні. Приклади гібридизації:

- CF + Content: при достатній кількості даних використовується колаборативна фільтрація, але для нових курсів або користувачів підмішується контентно-орієнтований рейтинг. Наприклад, система може спочатку відібрати список курсів за схожістю тем до вподобань студента, а потім відсортувати їх за колективним рейтингом інших користувачів. Приклади застосування такої гібридної моделі наведено у статті [2].

- Knowledge-based + CF: система враховує явні вимоги (рівень, пререквізити) і серед тих курсів, що підходять за цими вимогами, застосовує колаборативну фільтрацію для ранжування за популярністю чи оцінками.

- Комбінація кількох алгоритмів CF: наприклад, об'єднання результатів user-based і item-based CF, або традиційного CF та графового методу.

- Ensemble/Meta-рекомендації: кілька окремих рекомендацій (наприклад, для відео, для статей, для тестів) генерують пропозиції, а потім метарівень об'єднує їх (враховує контекст чи пріоритети користувача).

У контексті онлайн-курсів гібридні моделі показують вищу якість, оскільки беруть до уваги і подібність за змістом, і колективний досвід, і структурні обмеження. Дослідження підтверджують, що поєднання різноманітних алгоритмів підвищує ефективність рекомендацій навчальних матеріалів[3]. Гібридні методи також гнучкіші: система може динамічно змінювати вагу компонентів. Якщо студент новий - більше спиратися на контент та знання, якщо ж він активно навчається і лише відгуки - сильніше задіяти CF. Виклики гібридів є складність реалізації та пояснення, потреба узгоджувати різні джерела даних (наприклад, як об'єднати бал популярності з балом відповідності знанням). Незважаючи на це, більшість комерційних платформ (Coursera, Udemu тощо) використовують саме гібридні рекомендаційні системи - комбінуючи рейтинги, теги, прогрес користувача і навіть ручні рекомендації експертів.

Knowledge-based рекомендаційні системи використовують явні знання про предметну область, вимоги і цілі користувача для підбору контенту. На відміну від статистичних CF чи контентних методів, які вчать на даних, knowledge-based підхід спирається на експертно задані моделі знань: онтології, граfi знань, правила, пререквізити, профілі компетенцій. В освіті це надзвичайно важливо - курси мають передумови, логічну структуру, відповідність рівню підготовки. Система повинна врахувати «знання про знання»: що треба вивчити спочатку, які теми пов'язані, що саме хоче опанувати студент. Основна ідея - це рекомендувати доцільний навчальний ресурс, виходячи з поточного стану знань учня та його навчальної мети, замість просто того, що він колись вполював.

Приклад: студент зазначає, що його мета - вивчити машинне навчання. Knowledge-based система спочатку визначить, які базові компетенції потрібні (лінійна алгебра, програмування Python, основи статистики), і перевірить профіль учня чи його результати тестів, щоб зрозуміти, що вже опановано. Якщо, скажімо, виявиться, що студент знає Python, але не знайомий зі статистикою, система порекомендує курс з основ статистики як наступний крок. Такий підхід використовує «карти знань» або граfi предметних залежностей. У дослідженнях відзначається, що для коректної рекомендації курсів потрібно враховувати структуру знань користувача і залежності між курсами - без цього студенти часто не знають, який курс обрати, і не розуміють, які знання їм потрібні[3]. Knowledge-based система усуває цю проблему, явно моделюючи, що кому підходить: наприклад, онтологія може містити відношення «Курс В є продовженням курсу А», «Тема X

охоплюється в курсі Y», «Для теми Z потрібне розуміння X» і т.д. На основі цього система формує персональний навчальний маршрут з логічно пов'язаних кроків.

До цього класу належать і case-based reasoning системи, що підбирають навчальні ресурси, схожі на ті, що підходили користувачам з аналогічним профілем знань чи цілями (в деякому сенсі це теж “collaborative”, але по знанню, а не вподобаннях). Правила і експертні системи - ще один підвид: наприклад, якщо студент вивчив курс X і склав тест на >80%, то рекомендувати йому курс Y, інакше запропонувати спочатку курс Z (пререквізит). Такі системи були популярні в класичних інтелектуальних навчальних системах (ITS), де навчання розгалужується за сценаріями.

Перевагами підходів на основі знань є врахування педагогічної доцільності - система не порадить курс невідповідного рівня або без необхідної бази, можливість персоналізації під цілі - двом учням з різними цілями (скажімо, одному потрібно підтягнути матаналіз, іншому - навчитись Python) будуть рекомендовані зовсім різні траєкторії навіть при схожій історії, прозорість - такі рекомендації легше обґрунтувати («вам радимо X, бо ви ще не опанували Y, а це потрібно для Z»). Недоліками систем даного класу є трудомісткість побудови й оновлення знань (потрібні експерти для створення онтологій, графів компетенцій), масштабованість - правило, що працює в одному домені, не перенесеш автоматично в інший, ризик застарівання (навчальні програми змінюються). Крім того, якщо профіль користувача неточний (наприклад, студент неявно знає якусь тему, але це не зафіксовано) - рекомендація може недооцінити або переоцінити його і запропонувати неоптимальний маршрут.

На практиці knowledge-based методи часто інтегруються. Наприклад, великі МООС-платформи будують графи знань: вузли - курси або поняття, ребра - «курс A передуює B» або «курс C охоплює тему D». Це використовується разом зі статистичними даними: система спочатку фільтрує курси, що відповідають кар'єрній цілі користувача і доступні йому за пререквізитами, а далі ранжує їх за рейтингами чи успішністю (гібридний підхід).

Онлайн-курси містять матеріали у різних форматах - відеолекції, текстові статті, книги, слайди, аудіоподкасти, інтерактивні вправи, тести, тощо. Система має враховувати формат, оскільки поведінка користувача і методи аналізу контенту відрізняються залежно від типу медіа:

- Відео. Відеолекції - основний формат у МООС. Для рекомендацій відео важливі дані про поведінку при перегляді: час перегляду, чи користувач додивився його до кінця, чи робив паузи або переглядав повторно певні фрагменти. У дослідженні [7] доведено, що короткі відео (до 6 хвилин) значно ефективніші у залученні студентів. Ці сигнали можуть вказувати на зацікавленість або складність матеріалу. Алгоритми можуть враховувати, що студент переглянув 80% відео A, а відео B кинув на середині - значить, B не варто радити. Для аналізу змісту відео застосовують розшифровки (транскрипти) та розпізнавання мови: текстову транскрипцію лекції можна обробити як звичайний документ, виділивши ключові теми. Також використовуються метадані: назва, опис, теги до відео. Особливістю відео є тривалість - короткі ролики (до 5 хв) можна рекомендувати як додаткові приклади або пояснення, а довгі лекції - як основний матеріал. Система може збалансовувати рекомендації, чергуючи довгі курси з короткими підсумковими відео для закріплення. Виклик у рекомендації відео - велика вага користувацьких уподобань: платформи типу YouTube відомі своїми потужними відео-рекомендаторами, які враховують і контент, і колаборативні сигнали (оцінки, кліки). В освітньому середовищі слід додатково контролювати, щоб алгоритм не повів студента вбік від навчальної мети (наприклад, не почав радити розважальні відео).

- Текстові матеріали. До них належать статті, розділи підручників, навчальні кейси, веб-сторінки з теорією, слайди лекцій (з текстовими маркерами). Для текстів доступний широкий спектр методів NLP: від класичного TF-IDF для пошуку схожих документів до сучасних LLM-embedding для семантичного зіставлення. Рекомендація текстових матеріалів часто схожа на рекомендацію наукових статей: враховується тематична близькість, цитування (в академічних курсах може бути граф цитувань або посилань), популярність серед колег. Важливо врахувати складність тексту: наприклад, для школяра і для студента ВНЗ потрібні різні рівні викладу однієї теми. Тут можуть допомогти мовні моделі, що визначають рівень складності тексту чи його відповідність певній аудиторії. З поведінкових даних - можна аналізувати час читання (чи дочитав текст до кінця), виділення або коментарі. Якщо студент пропустив текст або читав дуже довго, система може запропонувати інший матеріал на ту ж тему, але простіший або детальніший. Текст також часто використовують як контекст для інших рекомендацій: наприклад, проаналізувавши відповіді студента в тесті (які теж текстові), система може порадити розділ теорії, де пояснюється незрозумілий концепт.

- Інші формати.

- Аудіо (подкасти, лекції) - схожі на відео без картинки: ключові параметри - тривалість, транскрипт, швидкість прослуховування. Можна рекомендувати аудіоматеріали для повторення в дорозі чи фонового навчання.

- Інтерактивні вправи і тести - їх рекомендують, щоб перевірити або закріпити знання. Тут рекомендація залежить від результатів: якщо студент зробив помилку в тесті з теми X, йому варто запропонувати ще завдання по X або повернутися до матеріалів із цієї теми. Рекомендаційна система може динамічно вставляти тестові запитання в навчальний шлях (adaptive assessment).

- Проекти та практичні завдання - рекомендуються на основі вже пройдених модулів. Наприклад, після кількох теоретичних курсів з програмування студенту можна порадити курсовий проект для практики. Алгоритми можуть використовувати профіль навичок: хто виконав проект з веб-розробки, тому буде цікавий наступний проект з близької галузі (скажімо, мобільна розробка).

○ Форуми та спільноти - нетиповий, але важливий ресурс. Деякі платформи радять обговорення або Q&A потоки, релевантні до курсу, який проходить студент. Тут в пригоді стає аналіз тексту питання/відповіді та соціальних вподобань (які треди читають чи лайкають схожі учні).

Успішна система повинна бути мультимодальною - тобто розуміти і відео, і тексти, і оцінки, та вміти комбінувати рекомендації різних типів. Наприклад, після перегляду відео про тему Y, система може порекомендувати тест для самоперевірки по Y, а потім статтю для поглиблення знань. Баланс форматів теж важливий: комусь зручніше читати, комусь дивитися - хороша рекомендація врахує індивідуальний стиль навчання користувача (якщо помітно, що він частіше обирає відео, система надасть пріоритет відеоматеріалам, і навпаки). Отже, різномірний контент збагачує освітню траєкторію, а завдання системи - обрати правильний формат в правильний момент.

Останні досягнення в галузі великих мовних моделей (GPT-3, GPT-4, PaLM, LLaMA тощо) відкривають нові можливості для персоналізації онлайн-навчання. LLM володіють потужними здібностями до розуміння природної мови, генерації тексту та семантичного узагальнення, що можна застосувати на різних етапах роботи системи планування навчання. Нижче розглянуті ключові напрямки, де інтеграція LLM здатна підвищити якість рекомендаційної системи:

Семантичний аналіз та анотування контенту. LLM можуть глибоко "читати" освітні матеріали - конспекти, статті, транскрипти лекцій і витягати з них змістовні характеристики. Замість ручної розмітки курсів по темах, модель на зразок GPT-4 здатна зрозуміти зміст уроку, виділити ключові поняття, визначити рівень складності, тон викладу. Це значно збагачує профіль контенту для контентно-орієнтованих і knowledge-based рекомендацій. Наприклад, LLM може прочитати новий курс і автоматично віднести його до певної галузі знань, знайти згадки пререквізитів (що треба знати перед цим курсом) та результатів навчання. Отримавши таку "семантичну анотацію", система краще зіставить курс з потребами студента. Крім того, LLM можна залучити для класифікації формату: визначити, що за тип матеріалу (лекція, практичне завдання, довідник) і відповідно вирішити, коли його доречно рекомендувати. Важлива перевага - LLM розуміють контекст і синоніми, тому різні формулювання тієї ж теми будуть розпізнані як пов'язані, чого не завжди досягти простими ключовими словами.

Моделювання знань і потреб користувача. Великі мовні моделі здатні аналізувати відкриті відповіді студентів, їх запити та навіть вести діалог зі студентом, щоб уточнити його цілі. Традиційно профіль учня складався з його успішності (бали, пройдені курси) або явних анкетних даних. З LLM з'являється можливість побудувати багатшу модель: наприклад, проаналізувати коротке есе або мотиваційний лист студента та витягти звідти інформацію про його інтереси, цілі навчання, попередній досвід. LLM можуть в реальному часі відповідати на запитання учня і по відповідях з'ясовувати, що йому потрібно краще пояснити. Концепція Knowledge Tracing (відстеження знань) теж еволюціонує з LLM: звичайні моделі знань (типу Bayesian Knowledge Tracing) оновлюють параметри знання за фактом правильності відповідей, а LLM може робити це ще й з огляду на пояснення відповіді. У нещодавній роботі TutorLLM запропоновано об'єднати класичний knowledge tracing з можливостями LLM: модель аналізує діалоги і відповіді студента, щоб оновлювати оцінку його знань і навичок [8]. Таке поєднання дозволяє врахувати, як студент мислить, де робить логічні помилки, і на цій основі персоналізувати рекомендації глибше, ніж просто правильно/неправильно. LLM перевершують традиційні системи в лінгвістичній гнучкості - користувач може запитати поради в довільній формі ("Хочу розібратись з алгоритмами індексування у реляційних базах даних. З чого мені почати?"), і модель здатна зрозуміти запит, розпитати уточнення і оновити модель потреб користувача. Таким чином, LLM додають елемент адаптивного репетитора, що слухає студента.

Персоналізована генерація навчальних траєкторій. Мабуть, найбільш інноваційне - доручити самій моделі генерувати поради щодо навчального плану. Сучасні LLM вже демонструють здатність створювати доволі осмислені навчальні плани на запит користувача [9]. Якщо дати моделі структуровану інформацію про доступні курси і знання студента, вона може згенерувати рекомендацію: "Спочатку пройдіть курс А (базовий), потім курс В для поглиблення, після чого виконайте проект С для практики". Дослідження показують, що при належному промптуванні (наданні інструкцій) LLM успішно будують персоналізовані плани навчання [9]. Наприклад, модель можна запитати: "Користувач знає основи статистики і хоче вивчити машинне навчання. У базі є курси: Linear Algebra, ML Basics, Deep Learning. Який шлях порадиш?". Модель, розуміючи логіку предметної області, може відповісти: "Спершу Linear Algebra (бо потрібна для ML), потім ML Basics". Такі рекомендації вже наближуються до рівня поради від досвідченого наставника. Перевагою є гнучкість: LLM можуть врахувати нестандартні вимоги ("вчу англійською, але хочу курс з українськими субтитрами") - модель підбере, якщо має цю інформацію). Реальні приклади: платформа Duolingo впровадила функції на основі GPT-4 для персоналізації навчання мові - модель генерує додаткові пояснення і практичні діалоги, пристосовані до рівня користувача. Отже, LLM можуть не тільки підбирати наявні матеріали, а й динамічно створювати контент: питання для самоперевірки, резюме пройденого, поради з тайм-менеджменту навчання тощо. Це збагачує персоналізацію: студент фактично отримує індивідуального асистента, інтегрованого з системою.

Покращення взаємодії та підтримки в навчанні. Окрім суто рекомендацій, LLM можуть підвищити якість спілкування системи з учнем. Замість статичного списку курсів, модель може пояснювати, чому рекомендується той чи інший крок, відповісти на уточнювальні запитання ("Чим корисний цей курс для моєї мети?"), підбадьорити або надати підказку при труднощах. Це важливий психологічний аспект: персоналізація полягає не лише в тому що рекомендувати, а й як подати рекомендацію. Наприклад, якщо студент кілька разів покидав курс на півдорозі, LLM-асистент може запропонувати інший формат навчання: "Бачу, що вам важко

даються довгі курси. Можливо, спробуємо серію коротких уроків з практикою?”. Така адаптивність в режимі діалогу значно перевершує традиційні системи, які зазвичай статично відображають рекомендації і не реагують на зворотний зв'язок в реальному часі [8]. Інший аспект - виявлення помилок і прогалин: аналізуючи відповіді студента на відкриті питання, LLM може з'ясувати, яке саме непорозуміння призвело до помилки, і порадити матеріал, що це пояснює. Це тонка персоналізація під конкретні прогалини в знаннях, яку складно досягти стандартними тестами.

Попри величезний потенціал, інтеграція LLM в освітні платформи має свої виклики. Один з них - галюцинації моделі, тобто ризик, що LLM згенерує впевнену, але невірну рекомендацію чи твердження. Тому все частіше використовується підхід Retrieval-Augmented Generation (RAG): модель отримує доступ до бази перевірених знань (описів курсів, статей) і зобов'язана спиратися на них при генерації відповіді. Це зменшує вигадки і робить рекомендації більш обґрунтованими фактами. Інший виклик - відсутність особистого досвіду: LLM не має емпатії чи реального педагогічного чуття, тому важливо обмежити її рамками, перевіряти відповіді. В ідеалі LLM-асистент має працювати під наглядом класичного алгоритму чи викладача: наприклад, пропонувати маршрут, а людина чи rule-based модуль затверджує, що він педагогічно коректний. Тим не менш, вже зараз LLM демонструють суттєві переваги у персоналізації навчання - від адаптивних підказок до повної побудови індивідуального курсу. Очікується, що подальша інтеграція LLM змінить освітні траєкторії: навчання стане більш діалоговим, гнучким і підлаштованим під кожного учня.

Нижче зведено основні типи рекомендаційних підходів, їх приклади застосування у персоналізації навчання, а також ключові сильні та слабкі сторони:

Таблиця 1

Підхід	Приклад	Сильні сторони	Недоліки
Колаборативна фільтрація	Рекомендація курсів, які обирали схожі студенти. Напр., студенту радять курс, що популярний серед інших з тієї ж спеціалізації (на основі історії виборів)	– Враховує колективний досвід – Не потребує аналізу змісту – Виявляє приховані інтереси (через уподобання групи)	– Cold start: нові користувачі/курси не мають даних – Може рекомендувати не по темі (тільки бо популярно) – Схильність до ефекту не персональних цілей
Контентно-орієнтований	Рекомендація матеріалу за схожим змістом. Напр., після курсу з алгоритмів радять книгу/статтю по алгоритмах (збігаються ключові слова, тема)	– Релевантність по темі і рівню – Працює для нових курсів (опис відомий) – Пояснюваність (спільні теми очевидні)	– Вузкість: радить дуже подібне (мало новизни); – Якість залежить від метаданих і NLP – Не враховує наскільки контент популярний чи цікавий іншим
Гібридна модель	Комбінація: фільтр спершу виключає курси без потрібних пререквізитів (knowledge-based), далі серед решти застосовує CF ранжування за рейтингом	– Поеднує переваги різних методів – Гнучкість під різні ситуації – Краще покриття сценаріїв (менше “сліпих зон”)	– Складність реалізації і налаштування – Важче пояснити користувачу (мікс факторів) – Потрібні дані одразу кількох видів (і метадані, і рейтинги, тощо)
Knowledge-based (на знаннях)	Порада навчального шляху за профілем знань. Напр., система знає пререквізити: після курсу “Calculus I” порадить “Calculus II”, а не одразу складний “Differential Eq.”	– Відповідність навчальним цілям і рівню – Уникнення недоречних рекомендацій (не пропустить важливу базу) – Прозорість (можна обґрунтувати правилами)	– Вимагає експертної моделі знань (онтології, графі) – Слабко вчиться з нових даних (правила треба оновлювати вручну) – Не враховує «смаків» користувача (тільки потреби) без додаткових модулів
Графові підходи	Рекомендація через зв'язки у графі. Напр., використання графа знань: студенту радять курс, який містить тему, пов'язану з уже пройденими темами, або GNN-модель пропонує курс, до якого багато коротких шляхів від його профілю	– Урахування складних взаємозв'язків (спільні теми, спільні студенти тощо) – Покращена персоналізація через графові embeddings – Може вирішувати cold start за рахунок знань (для нового курсу відомі зв'язки в графі)	– Трудомісткість та «чорний ящик» (важко пояснити рекомендацію через багаторівневі зв'язки) – Потребує актуального графа (знань або взаємодій), який треба підтримувати – Обчислювально інтенсивний для великих платформ

Насправді, межі між категоріями розмиті. Більшість реальних систем комбінують елементи всіх підходів. Наприклад, використання графа знань можна розглядати і як knowledge-based, і як графовий метод одночасно. Так само, LLM-інтегровані рішення часто є гібридними: модель генерує рекомендацію, але керується правилами (knowledge-based) і перевіряється на даних уподобань інших (CF). Тому при проектуванні сучасної рекомендаційної системи для онлайн-навчання зазвичай будують багатоваріантний підхід, де різні алгоритми доповнюють один одного на окремих етапах.

Аналіз сучасних підходів побудови персоналізованого навчального шляху

У роботі [10] представлено метод KG-PLPPM (Knowledge Graph-Based Personal Learning Path Planning Method), призначений для побудови персоналізованих навчальних траєкторій в онлайн-освіті. Підхід поєднує онтологічне моделювання предметної області з даними про поведінку студентів для підвищення ефективності навчання та усунення прогалин у знаннях. Для цього на основі відкритого набору даних MOOCube було побудовано граф знань, який відображає курси, поняття, відео та їхні зв'язки. Схожість між поняттями визначається двома способами: за семантичними ембеддингами TransR та за поведінковими даними через колаборативну фільтрацію. Обидві оцінки інтегруються у фінальний показник, який використовується разом із модифікованою когнітивною моделлю V-DINA, що оцінює рівень засвоєння понять за логами перегляду відео. На основі цих даних алгоритм формує індивідуальний шлях навчання, починаючи з найбільш слабких та релевантних понять, додаючи необхідні пререквізити у правильному порядку.

Експерименти на даних MOOCube показали, що KG-PLPPM суттєво перевершує базові методи Trans-CF та Ontology-CF. Метод забезпечує більш високу послідовність знань (0,29-0,60 проти $\approx 0,004$) та значне зростання ефективності навчання (покращення засвоєння від 104% до 148% залежно від сценарію). Це свідчить про здатність підходу формувати структуровані та когнітивно узгоджені траєкторії, що підвищують результативність навчання.

Щоб краще проілюструвати логіку роботи методу, на рис. 3 наведено схему загальної процедури KG-PLPPM: від збору та обробки даних до формування персоналізованої навчальної траєкторії.

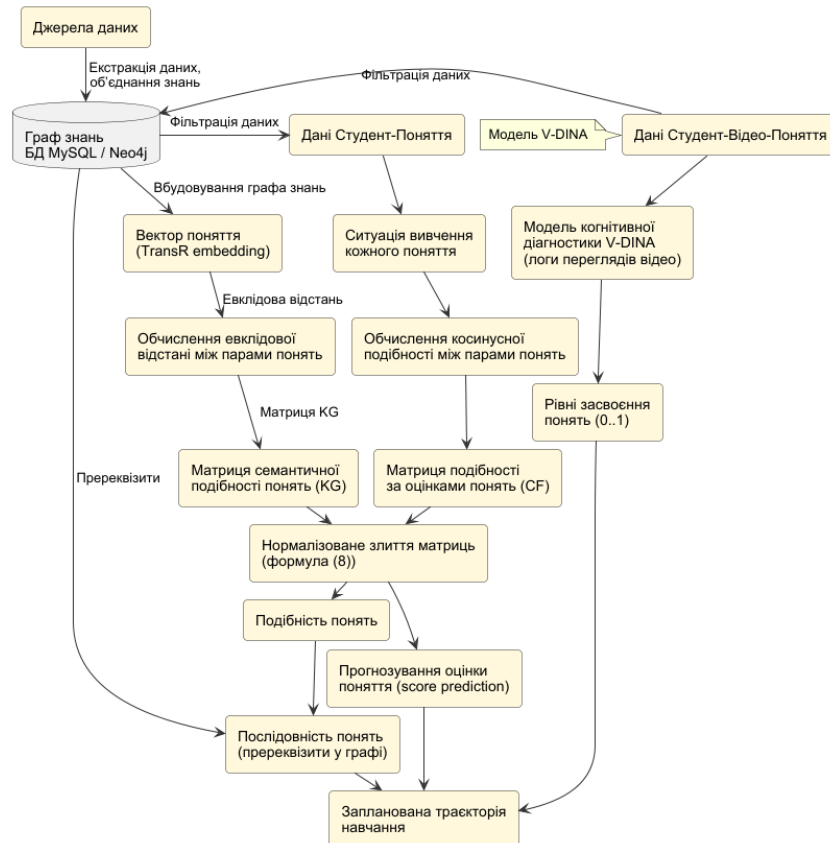


Рис. 3 Загальна процедура методу KG-PLPPM для планування персоналізованих навчальних траєкторій

Дане дослідження може слугувати міцною базою для порівняння та розвитку. Зокрема, можливе використання LLM і методів NLP для автоматизованого поповнення графа знань та уточнення зв'язків, розширення метрик схожості поняттями з текстових описів, а також застосування мультимодальних ознак (кліки, паузи, інтерактивність) для більш точної оцінки рівня засвоєння. Серед переваг методу — пояснюваність, поєднання семантики та поведінкових даних, фокус на слабких знаннях; серед обмежень — фіксовані параметри довжини шляху, залежність точності діагностики від обмежених поведінкових даних та відсутність природномовного інтерфейсу.

Продовжуючи графову лінію, однак відходячи від поєднання семантичних/поведінкових сигналів до збагачення самої структури знань, наступна робота фокусується на перетворенні ієрархії LOs у контекстуальний граф із новими семантичними ребрами [11]. Метою є подолання обмежень традиційної

ієрархії та відображення реального навчального контексту, що дозволяє пов'язувати LOs з різних програм та доменів за спільними цілями чи сценаріями застосування.

Методологія базується на кастомізованому конвеєрі text mining (TMP), який обробляє назви та описи LOs англійською та німецькою мовами. Для цього використано автоматичне визначення мови, очищення тексту, виявлення ключових тем (KeyBERT), формування багатомовних ембедингів (Sentence-BERT + Spacy, із перекладом DeepL) та обчислення косинусної подібності. Якщо подібність перевищує експериментально підібраний поріг, між об'єктами створюється новий тип зв'язку *has_semantic_relation_to*. Таким чином, до початкових експертно визначених зв'язків додаються автоматично згенеровані семантичні відношення, що формують щільніші контекстні підграфи.

Оцінювання виконувалося кількісно та якісно. Кількісна частина включала порівняння автоматично згенерованих зв'язків з експертними, а також аналіз мережових метрик. Показано, що 79% нових зв'язків мають подібність не нижчу за середній рівень. Перехід до KG підвищив середню ступеневу центральність (1,079 → 2,262), збільшив кількість спільнот (253 → 541) при зменшенні модульності, скоротив кількість слабо зв'язаних компонент (63 → 35) та більш ніж у дев'ять разів підвищив посередницьку центральність (1,57 → 15,1). Якісна оцінка з фокус-групами експертів у VET і TEL підтвердила, що додаткові семантичні зв'язки підвищують цінність рекомендацій, сприяють виявленню міждоменних “мостів” та підтримують багатомовність.

Даний підхід може слугувати як модуль семантичного збагачення графа знань. Його інтеграція дозволить автоматично доповнювати граф навчальних ресурсів прихованими контекстними зв'язками, покращувати рекомендації за рахунок урахування міждоменних відносин та підтримувати обробку природномовних запитів користувача з використанням LLM. Серед переваг методу — підвищення зв'язності та міждоменних переходів між LOs, підтримка багатомовності, використання семантичного аналізу для виявлення прихованих контекстів; серед обмежень — залежність від обсягу та якості текстових описів і потреба в адаптації моделей під конкретний домен.

Щоб оцінити, як графові представлення масштабуються до складних доменів і які моделі здатні уловлювати як локальні, так і глобальні залежності, звернімося до оглядової роботи з GNN — це дасть нам інструментарій, релевантний і для освітніх графів. У роботі [12] представлено огляд методів графового глибокого навчання, зосереджений на їх застосуванні для медичної діагностики та аналізу. Автори підкреслюють, що медичні дані часто мають неєвклідову природу (наприклад, зв'язки між клітинами, нейронами, симптомами чи зображеннями), тому класичні методи обробки даних (CNN, RNN) не здатні ефективно відобразити складні структурні залежності. Для цього пропонується використовувати графові нейронні мережі (GNN), які узагальнюють ідею згорткових операцій на графах і дозволяють враховувати як локальні патерни, так і глобальну структуру. Формально, оновлення представлення вузлів у GCN відбувається за правилом:

$$H^{(l+1)} = \sigma \left(\tilde{D}^{-\frac{1}{2}} \tilde{A} \tilde{D}^{-\frac{1}{2}} H^{(l)} W^{(l)} \right), \quad (4)$$

Де, $\tilde{A} = A + I$ - матриця суміжності з одиничною матрицею, $\tilde{D}$ - діагональна матриця степенів вузлів, $H^{(l)}$ - ознаки на шарі l , $W^{(l)}$ - ваги моделі, σ - функція активації.

Методологічно стаття систематизує архітектури GNN (Graph Convolutional Networks, Graph Attention Networks, GraphSAGE, Graph Isomorphism Networks тощо) та аналізує їх застосування до різних типів медичних даних: функціональної конективності мозку (fMRI, EEG), структурних зображень (MRI, CT, PET), біомаркерів і клінічних показників. Окрему увагу приділено інтеграції мультимодальних даних і побудові графів, де вузлами є пацієнти чи біологічні ознаки, а ребрами — зв'язки подібності або причинності. Такі підходи продемонстрували здатність підвищувати точність діагностики при складних захворюваннях (аутизм, хвороба Альцгеймера, Паркінсона, депресія, епілесія), а також допомагають виявляти приховані закономірності, недоступні для традиційних моделей.

Автори також окреслюють виклики: необхідність у великих високоякісних датасетах, висока обчислювальна складність, проблема узагальнення моделей між різними доменами, а також обмежена пояснюваність рішень GNN. Перспективними напрямками розвитку визначено створення динамічних графів, підвищення інтерпретованості результатів та ефективну інтеграцію графових моделей з іншими методами глибокого навчання.

Ця робота є важливим оглядом можливостей графових нейронних мереж, який демонструє їхній потенціал для моделювання складних структур і багатомодальних даних. У контексті персоналізованого планування навчальних траєкторій такий підхід може слугувати аналогією: студенти та навчальні ресурси можуть бути подані у вигляді графа з різними типами зв'язків (пререквізити, теми, результати тестів, стилі навчання). Перевагою є здатність GNN моделювати як локальні (зв'язки між близькими поняттями), так і глобальні (міждоменні переходи) відношення, що підвищує релевантність рекомендацій. Водночас обмеження залишаються схожими: потреба у великих масивах якісних освітніх даних, значні обчислювальні витрати та складність пояснення отриманих рекомендацій студентам і викладачам.

Переходимо від загального апарату графових мереж до конкретних освітніх систем, де ядром виступає когнітивна діагностика та експлуатація поведінкових логів у MOOCs для побудови динамічних траєкторій. У роботі [13] запропоновано динамічний алгоритм побудови персоналізованої навчальної траєкторії для середовища MOOCs, який враховує поточний стан засвоєння знань студентом та складність кожного знання.

На відміну від традиційних підходів, де складність маркується вручну, автори розробили модель оцінки складності знань на основі історичних даних про успішність, кількість повторних переглядів відео та коментарів (з використанням АНР для зважування). Паралельно створено модель засвоєння знань, яка діагностує стан студента (незасвоєно, не опановано, недостатньо опановано, опановано) на основі результатів тестів і поведінки при перегляді відео.

Алгоритм поєднує ці дві моделі та формує індивідуальний план, у якому прогалини в знаннях закриваються перед переходом до нових тем. Рекомендації генеруються у порядку пререквізитів та зростання складності. Для оцінки використано дані MOOCsCubeX та API платформи XuetangX. Експерименти показали, що група студентів, чия траєкторія збігалася з алгоритмічною, мала вищу ефективність навчальних дій (93% проти 71%), більший рівень завершення курсу (76% проти 57%) та вищі фінальні бали (75,2 проти 54,5) при меншому середньому часі навчання.

Данна робота є хорошим прикладом поєднання когнітивної діагностики з автоматичною оцінкою складності матеріалу. Цей підхід можна інтегрувати як ядро системи адаптивного планування, доповнивши методами NLP/LLM для автоматичного вилучення понять і пререквізитів, а також мультимодальними ознаками взаємодії (наприклад, участь у форумах, робота з інтерактивними вправами). Серед переваг — висока адаптивність, автоматизація оцінки складності, узгодженість із когнітивними моделями; серед обмежень — залежність від повноти даних логів, складність адаптації під офлайн-навчання та потреба у масштабуванні для різних предметних областей.

Наступний крок — уніфікувати діагностику знань із добором послідовності і ресурсів: мультиалгоритмічний підхід поєднує Fuzzy-CDF для «прогалин», Аргіорі для порядку понять та PSO для підбору матеріалів за профілем студента[14]. Автори виходять з того, що надмірність ресурсів у середовищі MOOCs призводить до інформаційного перевантаження та втрати цілісності навчання, і пропонують системний підхід, що враховує чотири аспекти профілю студента: когнітивний рівень, навчальну здатність, стиль навчання та інтенсивність навчання.

Методологія передбачає два основні етапи. Спочатку на основі моделі Fuzzy-CDF оцінюється ступінь засвоєння знань і визначається набір “незасвоєних” понять. Далі з використанням алгоритму аргіорі та карти знань генерується оптимальна послідовність цих понять, враховуючи їх пререквізити та взаємозв'язки. На другому етапі до кожного знання підбираються найбільш релевантні ресурси за допомогою PSO, з урахуванням відповідності типу ресурсу стилю навчання студента, складності матеріалу його навчальній здатності та часу освоєння — інтенсивності навчання.

Експерименти на даних з 584 математичних понять середньої школи показали, що така комбінація алгоритмів забезпечує покриття понад 90% необхідних знань та підвищує точність рекомендацій порівняно з топологічним сортуванням, алгоритмом світлячків, методами трансферного навчання та класичними DL-моделями. Крім того, студенти, які навчалися за запропонованими траєкторіями, продемонстрували вищі підсумкові бали при меншому часі навчання.

Модель, описана в цій роботі, цікава як приклад інтеграції кількох методів ШІ для побудови гнучких навчальних маршрутів, особливо у поєднанні когнітивної діагностики та оптимізаційних алгоритмів для підбору ресурсів. Переваги підходу - багатофакторне моделювання студента, адаптація ресурсів до індивідуальних особливостей та стійка робота PSO. Недоліки - висока складність побудови детального профілю студента та потреба у великому обсязі логів взаємодії для якісної діагностики.

Щоб додати вимір індивідуальних стилів навчання і персональних уподобань у виборі активностей, наступна робота переходить до гібриду fuzzy-ANN з подальшою агрегацією WSM [15]. У роботі представлено гібридну модель fuzzy-ANN для персоналізованих рекомендацій навчальних активностей, побудовану на основі теорії експериментального навчання Колба (Kolb's Experiential Learning Theory). Автори підкреслюють, що класичні rule-based підходи або окреме використання нечіткої логіки та нейронних мереж мають обмеження у відображенні індивідуальних стилів навчання. Запропонована модель поєднує переваги нечітких систем, які дозволяють формалізувати невизначеність у визначенні стилю, та здатність ШНМ вчитися з даних і враховувати попередню історію взаємодії студента.

Методологія включає два основні етапи. Спершу індивідуальні стилі студентів визначаються за допомогою опитувальника Kolb's Learning Style Inventory (LSI), результати якого трансформуються у набір нечітких функцій належності (ступені приналежності до стилів Accommodator, Converger, Assimilator, Diverger). Ці fuzzy-вектори разом із даними про історію навчання (результати виконання завдань, взаємодію з системою) подаються на вхід штучної нейронної мережі, яка формує рекомендації шести типів навчальних активностей: практичні (hands-on), проблемно-орієнтовані, теоретичні, рефлексивні, колаборативні та з покроковим супроводом (guided). На другому етапі результати уточнюються за допомогою методу Weighted Sum Model (WSM), який агрегує ефективність активностей на основі фідбеку студентів (рейтинги від 1 до 5) та результатів схожих профілів, що підвищує узгодженість рекомендацій із реальними вподобаннями та результативністю навчання.

Експерименти проводилися на курсі з програмування C++ тривалістю 15 тижнів за участі 100 студентів. Запропонована fuzzy-ANN модель досягла точності 88,4%, перевищивши показники трьох базових моделей (rule-based, fuzzy-only, ANN-only). За метриками Precision (85,2%), Recall (87,8%) та F1-score (86,8%) вона також продемонструвала найкращі результати. Крім кількісної оцінки, проведено якісне case study для трьох

студентів, що показало здатність моделі індивідуалізувати траєкторію навіть для учнів зі схожими стилями, але різними патернами поведінки.

Дана робота є прикладом інтеграції когнітивної теорії навчання з методами штучного інтелекту для створення персоналізованих рекомендацій. У контексті моєї теми цей підхід демонструє, як можна поєднати структуровану модель стилів навчання з адаптивними можливостями нейронних мереж. Перевагами є багатозадачне представлення індивідуального профілю, врахування не тільки початкового стилю, але й динамічного фідбеку, а також підвищення точності рекомендацій завдяки гібридності. Серед обмежень - потреба у масштабованих і якісних даних анкетування, складність формування достатньо великої навчальної вибірки для ШНМ та ризик перевантаження системи при великій кількості активностей.

У роботі [16] представлено метод побудови персоналізованих навчальних траєкторій на основі багатозадачного глибинного навчання з використанням LSTM. Автори розглядають задачу рекомендації навчального шляху як задачу Seq2Seq-прогнозування, де вхід - це історія взаємодії студента з навчальними елементами, а вихід - послідовність рекомендованих кроків навчання. Модель одночасно вирішує дві задачі: (1) прогнозування наступних навчальних концептів (learning path recommendation) та (2) глибоке трасування знань (deep knowledge tracing) для оцінки ймовірності успішного засвоєння рекомендованого шляху.

Архітектура складається з спільного шару LSTM для вилучення загальних ознак обох задач та двох окремих LSTM-шарів, спеціалізованих під кожну з них. Додатково використано non-repeat loss, що штрафувє повтори концептів у траєкторії, стимулюючи різноманітність і логічну послідовність. Модель тренувалася на датасеті ASSIST09 з історіями взаємодії студентів, де тестували різні довжини навчальних шляхів (3, 5, 7, 9 елементів).

Загальна функція втрат моделі об'єднує обидва завдання та штраф за повтори у вигляді

$$L_{total} = L_{CE} + \lambda_1 L_{BCE} + \lambda_2 L_{rep} \quad (5)$$

що дозволяє одночасно оптимізувати якість рекомендацій, точність трасування знань та різноманітність шляху.

Результати показали стабільну перевагу над шістьма базовими моделями (RNN, LSTM, Seq2Seq з/без Attention): досягнуто Accuracy = 0,3489, F1-score = 0,3241, що суттєво вище за найближчі аналоги. Експерименти також показали, що ефективність падає зі збільшенням довжини шляху, але модель зберігає перевагу над конкурентами навіть для довших послідовностей.

Дана робота цікава як приклад інтеграції прогнозування знань та рекомендацій у спільному навчанні, що дозволяє покращити узагальнення та адаптивність моделі. Переваги - підвищена точність завдяки багатозадачності, зменшення перенавчання, забезпечення різноманітності шляху. Обмеження - орієнтація на послідовні логи даних (що потребує деталізованих записів взаємодій) та відсутність явної роботи з семантикою знань, яку можна було б посилити за допомогою графів знань і LLM.

У роботі [17] представлено метод персоналізованого планування навчальних траєкторій (PLPP) із використанням великих мовних моделей (LLM) та prompt engineering для підвищення адаптивності, інтерактивності та прозорості рекомендацій. Автори використовують моделі Llama 2 (70B) та GPT-4, формуючи серію спеціально сконструйованих підказок, які враховують індивідуальні характеристики студента - попередні знання, цілі навчання та проблемні зони. Система підтримує багатотурові діалоги, уточнюючи потреби користувача, та вбудовує пояснення до рекомендацій, що підвищує довіру й розуміння логіки траєкторії.

Метод включає три основні типи підказок: (1) Initial Assessment Prompt - визначає вихідний рівень і пропонує наступні концепти. (2) Clarification Prompt - уточнює складні для студента аспекти. (3) Explanatory Prompt - пояснює педагогічну доцільність порядку вивчення. Оцінка проводилася за метриками точність, задоволеність користувачів і якість траєкторій. Результати показали значне перевищення базового підходу: для GPT-4 точність зросла з 75,8% до 88,3%, задоволеність - з 3,7 до 4,4, а якість - з 3,9 до 4,7. Додатковий тримісячний аналіз підтвердив підвищення тестових балів на 21,4% і зростання рівня утримання студентів до 82,5%.

Даний підхід цікавий як інтерактивний модуль генерації навчальних шляхів на основі LLM з вбудованими поясненнями та адаптивним уточненням. Це дозволяє інтегрувати природномовний інтерфейс для запитів і підвищити прозорість рекомендацій. Переваги - висока адаптивність, підвищення довіри завдяки поясненням, краща якість траєкторій. Недоліки - залежність від якості вхідних даних, обчислювальна вартість роботи з великими LLM та потреба в ретельному проєктуванні prompt-структур.

Стаття [18] пропонує багатомодальний підхід до планування персоналізованих навчальних траєкторій у вищій освіті, що поєднує глибинні генеративні моделі, трансформери, змагальне (adversarial) навчання та квантову класифікацію станів. Система інтегрує текстові, аудіо- та відеодані студентів, використовуючи BERT, STFT/MFCC та ResNet для вилучення ознак, вирівнює їх у єдиному вбудованому просторі та виконує глибинне злиття через варіаційний автокодер. Трансформер з механізмом самоуваги аналізує динамічні зміни інтересів і статусу навчання, забезпечуючи адаптивну генерацію траєкторій у реальному часі. Для підвищення стійкості до аномальних даних використовується змагальне навчання з генерацією тестових збурень, а квантова класифікація прискорює обробку багатовимірних ознак, скорочуючи час обчислень із 87 до 56 секунд на 18-й ітерації.

На рис. 4 зображено загальну схему персоналізованого планування навчального шляху, запропоновану у даній роботі:

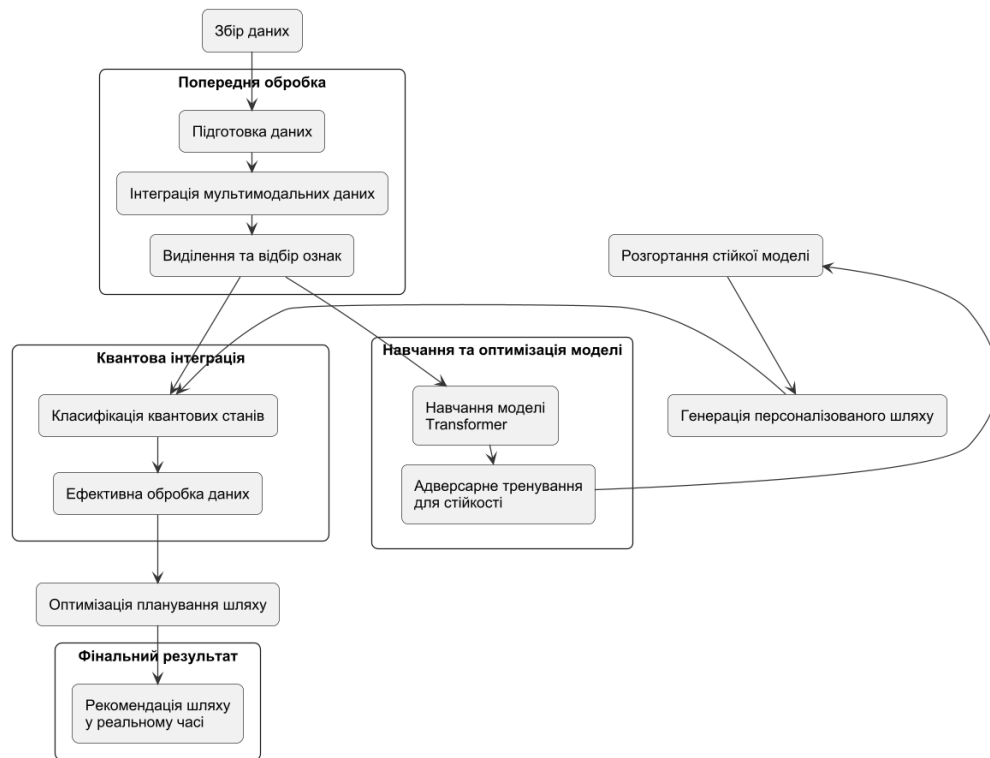


Рис. 4 Схема персоналізованого планування навчального шляху

У порівнянні з колаборативною фільтрацією та LSTM-моделлю, запропонований метод досягнув точності 0,95 і повноти 0,91, перевищивши аналоги. Перевагою підходу є здатність комплексно інтегрувати різномодальні дані та реагувати на зміни в навчанні. Недоліками є висока обчислювальна складність трансформерів та квантових методів, потреба у складному апаратному забезпеченні та труднощі забезпечення інтерпретованості рішень.

Висновки

У даній статті проведено систематизований і критичний аналіз сучасних підходів до персоналізованого планування навчальних траєкторій, що дало змогу узагальнити основні напрями розвитку в цій сфері та виокремити ключові методологічні групи. Розглянуті дослідження показали, що графові моделі знань (у тому числі контекстуальні KG) забезпечують формальне представлення предметних областей і пререквізитів, створюючи основу для побудови когнітивно узгоджених шляхів. Методи когнітивної діагностики дозволяють оцінювати індивідуальний рівень засвоєння понять і виявляти прогалини, тоді як комбінаторно-оптимізаційні схеми (наприклад, Fuzzy-CDF + Apriori + PSO) додають механізми автоматичного впорядкування знань і добору ресурсів. Нечітко-нейронні моделі, побудовані на основі стилів навчання, відкривають можливості для індивідуалізації не лише контенту, а й форм подання матеріалу. Послідовні та багатозадачні архітектури (Seq2Seq, LSTM) демонструють ефективність у прогнозуванні подальших кроків навчання, враховуючи часову динаміку даних і закономірності поведінки студентів. Мультимодальні рішення та великі мовні моделі (LLM) розширюють горизонти персоналізації завдяки інтеграції різних типів даних (текстових, відео, поведінкових) і створюють умови для природномовної взаємодії з освітніми системами.

Здійснений огляд підтвердив, що кожен із підходів має свої переваги та обмеження. Так, графові методи відзначаються пояснюваністю та структурованістю, проте вимагають трудомісткої побудови й підтримки онтологій. Когнітивна діагностика дає точне уявлення про рівень опанування знань, але потребує повних і якісних логів взаємодії. Комбінаторні й оптимізаційні алгоритми добре працюють із формалізованими даними, проте їх продуктивність обмежена масштабом і складністю вхідних моделей. Нечітко-нейронні системи зі стилями навчання дозволяють враховувати індивідуальні відмінності, однак потребують анкетування й значної кількості якісних прикладів для тренування. Послідовні й багатозадачні нейромережі здатні моделювати динаміку засвоєння знань, проте залишаються вимогливими до обчислювальних ресурсів та обсягів даних. Мультимодальні підходи й LLM вносять гнучкість і природність взаємодії, але водночас створюють нові виклики, пов'язані з валідацією результатів, контролем якості та інтерпретованістю.

Отримані результати свідчать про необхідність інтеграції цих підходів у гібридні архітектури, які б поєднували структурованість графів знань, точність когнітивної діагностики, адаптивність нейромереж і гнучкість LLM. Подальші дослідження доцільно зосередити на розробці комплексних моделей, здатних працювати з обмеженими або частково заповненими даними, забезпечувати баланс між точністю рекомендацій та прозорістю їх пояснень, а також масштабуватися до великих освітніх середовищ на кшталт MOOCs. Особливого значення набувають питання інтеграції мультимодальних сигналів, уніфікації метрик оцінювання якості навчальних траєкторій і створення протоколів відтворюваності результатів. Узагальнюючи, можна стверджувати, що перспективним напрямом є розробка інформаційних технологій, які поєднують когнітивну

діагностику, графові подання знань, оптимізаційні методи та LLM, що у сукупності дозволить забезпечити масштабоване, адаптивне й пояснюване персоналізоване навчання.

Література

1. Jin, Shasha, and Lei Zhang. "Research on the Recommendation System of Music E-Learning Resources with Blockchain Based on Hybrid Deep Learning Model." *Scalable Computing: Practice and Experience*, vol. 25, no. 3, 12 Apr. 2024, pp. 1455-1465, <https://doi.org/10.12694/scpe.v25i3.2672>
2. O. Kopylchak, and I. Kazymyра. "Гібридна модель для ефективного формування персоналізованих рекомендацій навчальних курсів." *COMPUTER-INTEGRATED TECHNOLOGIES EDUCATION SCIENCE PRODUCTION*, vol. 57, no. 57, 13 Feb. 2025, pp. 91-100, <https://doi.org/10.36910/6775-2524-0560-2024-57-11>
3. Lu, Jinliang. "A Survey of Online Course Recommendation Techniques." *Open Journal of Applied Sciences*, vol. 12, no. 01, 2022, pp. 134-154, <https://doi.org/10.4236/ojapps.2022.121010>
4. Gao, Yunmei. "MOOCs Video Recommendation Using Low-Rank and Sparse Matrix Factorization with Inter-Entity Relations and Intra-Entity Affinity Information." *Information Processing & Management*, vol. 61, no. 6, 6 Aug. 2024, p. 103861, [www.sciencedirect.com/science/article/pii/S0306457324002206, https://doi.org/10.1016/j.ipm.2024.103861](https://doi.org/10.1016/j.ipm.2024.103861)
5. Oliveira, Alana, et al. "Recommendation of Educational Content to Improve Student Performance: An Approach Based on Learning Styles." *Proceedings of the 12th International Conference on Computer Supported Education*, 2020, <https://doi.org/10.5220/0009436303590365>
6. Li, Bifeng, et al. "A Personalized Recommendation Framework Based on MOOC System Integrating Deep Learning and Big Data." *Computers and Electrical Engineering*, vol. 106, 1 Mar. 2023, p. 108571, [www.sciencedirect.com/science/article/pii/S0045790622007868?via%3Dihub, https://doi.org/10.1016/j.compeleceng.2022.108571](https://doi.org/10.1016/j.compeleceng.2022.108571)
7. Guo, Philip & Kim, Juho & Rubin, Rob. (2014). How video production affects student engagement: An empirical study of MOOC videos. 41-50. 10.1145/2556325.2566239. <http://dx.doi.org/10.1145/2556325.2566239>
8. Li, Z., Yazdanpanah, V., Wang, J., Gu, W., Shi, L., Cristea, A. I., Kiden, S., & Stein, S. (2025). TutorLLM: Customizing Learning Recommendations with Knowledge Tracing and Retrieval-Augmented Generation. *ArXiv.org*. <https://arxiv.org/abs/2502.15709>
9. Wang, X. J., Lee, C. P., & Mutlu, B. (2025). LearnMate: Enhancing Online Education with LLM-Powered Personalized Learning Plans and Support. 1-10. <https://doi.org/10.1145/3706599.3719857>
10. Hou, B., Lin, Y., Li, Y., Fang, C., Li, C., & Wang, X. (2025). KG-PLPPM: A Knowledge Graph-Based Personal Learning Path Planning Method Used in Online Learning. *Electronics*, 14(2), 255-255. <https://doi.org/10.3390/electronics14020255>
11. Hasan Abu-Rasheed, Mareike Dornhöfer, Weber, C., Gábor Kismihók, Buchmann, U., & Fathi, M. (2023). Building Contextual Knowledge Graphs for Personalized Learning Recommendations Using Text Mining and Semantic Graph Completion. *ArXiv (Cornell University)*, 36-40. <https://doi.org/10.1109/icalt58122.2023.00016>
12. Ahméd-Aristizabal, D., Armin, M. A., Denman, S., Fookes, C., & Petersson, L. (2021). Graph-Based Deep Learning for Medical Diagnosis and Analysis: Past, Present and Future. *Sensors*, 21(14), 4758. <https://doi.org/10.3390/s21144758>
13. Jiang, B., Li, X., Yang, S., Kong, Y., Cheng, W., Hao, C., & Lin, Q. (2022). Data-Driven Personalized Learning Path Planning Based on Cognitive Diagnostic Assessments in MOOCs. *Applied Sciences*, 12(8), 3982. <https://doi.org/10.3390/app12083982>
14. Ma, Y., Wang, L., Zhang, J., Liu, F., & Jiang, Q. (2023). A Personalized Learning Path Recommendation Method Incorporating Multi-Algorithm. *Applied Sciences*, 13(10), 5946. <https://doi.org/10.3390/app13105946>
15. Christos Troussas, Akrivi Krouska, Mylonas, P., & Sgouropoulou, C. (2025). A Fuzzy-Neural Model for Personalized Learning Recommendations Grounded in Experiential Learning Theory. *Information*, 16(5), 339. <https://doi.org/10.3390/info16050339>
16. Nasrin, A., Qian, L., Obiomon, P., & Dong, X. (2025). Enhancing Learning Path Recommendation via Multi-task Learning. *ArXiv.org*. <https://arxiv.org/abs/2507.05295>
17. Sun, C., Huang, S., Sun, B., & Chu, S. (2025). Personalized learning path planning for higher education based on deep generative models and quantum machine learning: a multimodal learning analysis method integrating transformer, adversarial training and quantum state classification. *Discover Artificial Intelligence*, 5(1). <https://doi.org/10.1007/s44163-025-00252-6>
18. Ng, C., & Fung, Y. (2024b). Educational Personalized Learning Path Planning with Large Language Models. *ArXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2407.11773>
19. Kopylchak, O., Kazymyра, I., Mukan, O., & Bondar, B. (2024). Personalized education plan construction using neural networks. *Mathematical Modeling and Computing*, 11(4), 1003-1012. <https://doi.org/10.23939/mmc2024.04.1003>