
THEORY INFORMATION TECHNOLOGIES AND SYSTEMS CONSTRUCTION

ТЕОРІЯ ПОБУДОВИ ІНФОРМАЦІЙНИХ ТЕХНОЛОГІЙ ТА СИСТЕМ

<https://doi.org/10.15407/intechsys.2026.03.003>

УДК 004.8:37

Ю.О. МАКОГОН, аспірант,
Харківський національний університет радіоелектроніки,
просп. Науки, 14, Харків, 61166, Україна
<https://orcid.org/0009-0005-4369-1555>
yurii.makohon@nure.ua

О.Г. РУДЕНКО, д-р. техн. наук, професор,
Харківський національний університет радіоелектроніки,
просп. Науки, 14, Харків, 61166, Україна
<https://orcid.org/0000-0003-0859-2015>
oleh.rudenko@nure.ua

АРХІТЕКТУРНИЙ СИНТЕЗ КОМПОЗИТНИХ НЕЙРОМЕРЕЖЕВИХ СИСТЕМ ДЛЯ ПРЕДИКТИВНОГО ОБРОБЛЕННЯ ДАНИХ У ЦИФРОВИХ ОСВІТНІХ СЕРЕДОВИЩАХ

Це дослідження обґрунтовує методологічні та технічні принципи синтезу композитних нейромережових архітектур, призначених для предиктивного оброблення даних у межах інфокомунікаційних освітніх середовищ. У контексті глобальної цифрової трансформації дослідження зосереджено на розробці адаптивних систем управління, де інтеграція рекурентних (RNN), графових (GNN) та трансформерних (Transformers) структур забезпечує високоточне моделювання когнітивної динаміки. У роботі розкрито функційні механізми динамічної оркестрації контенту та предиктивного моніторингу патернів оброблення інформації як основних елементів стратегії оптимізації. Проаналізовано технічні та епістемологічні обмеження, зокрема алгоритмічну непрозорість та інформаційну асиметрію, що зумовлює необхідність переходу до гібридних моделей для підвищення інтерпретованості інтелектуальних рішень. Особливу увагу приділено стратегічній ролі оптимізованих систем штучного інтелекту (ШІ) в Україні як життєво важливого інструменту для подолання освітніх розривів у кризові періоди. Переваги композитного підходу експериментально доведено на апаратному забезпеченні з обмеженими ресурсами (1 ГБ відеопам'яті), де було досягнуто значного скорочення затримки та підвищення стабільності системи.

Цитування: Макогон Ю.О., Руденко О.Г. Архітектурний синтез композитних нейромережових систем для предиктивного оброблення даних у цифрових освітніх середовищах. *Information Technologies and Systems*. 3 (9). 2026. 3 – 14. <https://doi.org/10.15407/intechsys.2026.03.003>

© Publisher PH “Akadempriodyka” of the NAS of Ukraine, 2025. This is an Open Access article under the CC BY-NC-ND 4.0 license (<https://creativecommons.org/licenses/by-nc-nd/4.0/>)

Ключові слова: адаптивне навчання, композитні нейромережеві архітектури, інфокомунікаційні системи, предиктивне оброблення даних, персоналізація, локальні обчислення, штучний інтелект.

Вступ

Поточна еволюція освітніх просторів зумовлена експоненційною інтеграцією цифрових інструментів, що призводить до виникнення складних гібридних інфокомунікаційних середовищ. Нейронні мережі (НМ) виступають основними технологічними агентами в цій трансформації, здатними виконувати сучасну предиктивну аналітику та системну оптимізацію траєкторій навчання. Впровадження інтелектуальних алгоритмів підкреслює критичну дихотомію між операційною ефективністю обчислювальних моделей та необхідністю збереження когнітивної автономії. Системні виклики охоплюють забезпечення алгоритмічної прозорості (*explainability*), мінімізацію інформаційної асиметрії та запобігання реплікації соціальних упереджень під час машинного навчання на нерепрезентативних наборах даних. Враховуючи, що композитні НМ оперують нелінійною динамікою, епістемологічні питання щодо природи знань, генерованих алгоритмами, стають дедалі значущими. В українському контексті використання ШІ для оптимізації засвоєння знань має стратегічний вимір, виступаючи ресурсом для підтримки когнітивної стійкості та подолання дефіциту компетентностей, спричиненого воєнним конфліктом. Відповідно, існує об'єктивна потреба в науковій базі для синтезу архітектур НМ з урахуванням соціотехнічних наслідків.

Попри значні результати, синтез композитних рішень для цілісної оптимізації траєкторій в українському контексті залишається недостатньо дослідженим. Відсутність методологій, що поєднують структурні можливості *GNN* із семантичною потужністю Трансформерів у межах однієї адаптивної моделі, визначає актуальність цієї роботи.

Метою статті є наукове обґрунтування методологічних принципів синтезу композитних архітектур НМ для предиктивного оброблення даних у межах цифрової трансформації. Дослідження зосереджено на розробці стратегії синергетичного поєднання *RNN*, *GNN* та Трансформерів для інтенсифікації адаптивного навчання та формалізації процедур предиктивного коригування контенту.

Аналіз та формальна постановка задачі оптимізації

У межах сучасної дидактичної парадигми освітній процес інтерпретується як нелінійна динамічна система, де результуюча ефективність визначається складною взаємодією з багатофакторним середовищем [1, 2]. Суб'єкт навчання в цій моделі виступає як активний когнітивний агент, чий внутрішній стан безперервно еволюціонує

під впливом зовнішніх керуючих впливів. Впровадження інтелектуальних систем дозволяє здійснювати предиктивне управління середовищем, спрямоване на досягнення екстремумів цільових функцій успішності.

Математичне моделювання коригування траєкторії виконується за допомогою рівняння стану (1):

$$S_{t+1} = f(S_t, A_t, E_t), \quad (1)$$

де S_{t+1} — прогнозований стан когнітивної системи учня у наступний момент часу; S_t — поточний стан когнітивної системи учня у момент часу t ; A_t — множина дій (навчальних завдань, взаємодій із вчителем чи системою); E_t — стохастичний вектор впливу зовнішнього середовища; f — функція переходу, що апроксимується нейронною мережею.

Використання композитних моделей забезпечує перехід від строго детермінованих схем до стохастичної адаптивної функції f , де параметри динамічно перераховуються на основі даних зворотного зв'язку. Цей підхід синтезує обчислювальний аналіз з оцінкою епістемологічних зрушень у структурі комунікації [3].

Для формальної верифікації якості оптимізації в *Learning Analytics* [4] використовуються цільові функції. Задачі класифікації рівнів засвоєння модулів вирішуються шляхом мінімізації втрат крос-ентропії (2):

$$L(y, \hat{y}) = - \sum y_i \log(\hat{y}_i), \quad (2)$$

де L — значення функції втрат; y — приналежність до реального класу (рівня знань); $\hat{y}$ — прогнозована ймовірність приналежності до класу.

Прогнозування безперервних показників (наприклад, ймовірності деградації знань або часових витрат) базується на мінімізації середньоквадратичної помилки (MSE) (3):

$$MSE = \frac{1}{n} \sum (y_i - \hat{y}_i)^2, \quad (3)$$

де n — кількість спостережень у вибірці; y_i — фактичне значення показника (наприклад, реальний бал учня); $\hat{y}_i$ — прогнозоване значення, отримане на виході нейронної мережі.

Аналіз нейромережових парадигм у системній оптимізації

Аналіз останніх розробок [4–6] свідчить про те, що вибір конкретної архітектури НМ є результатом дидактичного проєктування, де параметри корелюють із характеристиками даних та цілями. RNN, GNN та Трансформери оцінюються у трьох доменах: темпоральному моделюванні, структурній організації та семантичній інтерпретації.

Темпоральне моделювання та предиктивна аналітика. Динамічна природа освіти вимагає врахування темпорального контексту. Математичні задачі містять апроксимацію ймовірності переходу (4):

$$P(S_t | S_{t-1}, A_{t-1}), \quad (4)$$

де P — умовна ймовірність переходу в новий стан знань; S_t — стан знань у поточний момент; S_{t-1} — стан знань на попередньому кроці; A_{t-1} — дія або навчальне завдання, виконане на попередньому кроці.

Long Short-Term Memory (LSTM) архітектури є класичними для цих завдань [7], забезпечуючи вектори прихованого стану для довготривалої пам'яті. Трансформери пропонують альтернативу, аналізуючи масиви часових рядів паралельно за допомогою позиційного кодування, уникаючи втрати контексту в довгих послідовностях.

Структурне моделювання когнітивних зв'язків. Когнітивні домени мають складні нелінійні топології. У той час як *RNN* мають труднощі з лінеаризацією графових даних, *GNN* нативно спроектовані для реляційних даних. Відповідно до праць [8, 9], середовища подають як формальні графи (5):

$$G = (V, E), \quad (5)$$

де V — множина вершин, що подають навчальні концепти або студентів; E — множина ребер, що відображають когнітивні або соціальні зв'язки між ними.

Механізм *Message Passing* забезпечує динамічне оновлення стану (6):

$$h_v^{(l+1)} = \sigma \left(\sum_{u \in N(v)} \frac{1}{c_{vu}} W^{(l)} h_u^{(l)} \right), \quad (6)$$

де $h_v^{(l+1)}$ — векторне подання вершини v на шарі $l + 1$; σ — нелінійна функція активації (наприклад, *ReLU*); $N(v)$ — множина сусідніх вершин для вершини v ; c_{vu} — нормалізуюча константа; $W^{(l)}$ — матриця ваг, що навчається; $h_u^{(l)}$ — стан сусідньої вершини u на попередньому шарі.

Семантична інтерпретація контенту. Оптимізація контенту вимагає контекстно-залежного оброблення природної мови (*NLP*). Трансформери [10], використовуючи *Self-Attention*, динамічно оцінюють значущість елементів незалежно від відстані. Це формалізується через *Scaled Dot-Product Attention* (7):

$$\text{Attention}(Q, K, V) = \text{soft max} \left(\frac{QK^t}{\sqrt{d_k}} \right) V, \quad (7)$$

де Q — матриця запитів, що представляють поточні завдання або інтереси учня; K — матриця ключів, що кодують характеристики навчального матеріалу або попередній досвід; V — матриця значень, що

містить змістову інформацію про контент; d_k — розмірність векторів ключів (використовується для масштабування градієнтів); *softmax* — функція активації для отримання розподілу ймовірностей уваги.

Це дає змогу вирішувати складні лінгвістичні проблеми, такі як полісемія та кореференція [11].

Апаратна інтенсивність та обчислювальна складність. Впровадження стратегій оптимізації лімітується наявним обладнанням. *RNN* демонструють лінійну складність $O(N)$, роблячи їх ефективними для локального та мобільного розгортання. Складність *GNN* залежить від щільності графа $O(\|V\| + \|E\|)$. Трансформери демонструють найвищу інтенсивність через квадратичну залежність $O(N^2)$, що вимагає спеціалізованих *GPU* або *TPU* кластерів. Порівняльні результати узагальнено в Таблиці 1.

Результати цього всебічного аналізу демонструють, що жодна монолітна архітектура не може самостійно задовольнити багатопланові вимоги інтелектуальної освітньої екосистеми. Ефективна оптимізація системи вимагає переходу до композитних ансамблів нейронних мереж, використовуючи специфічні обчислювальні та когнітивні переваги кожної парадигми для балансу точності та ресурсоефективності.

Експериментальна оцінка продуктивності в умовах жорстких апаратних обмежень

Для підтвердження практичної спроможності синтезованих композитних архітектур було проведено серію тестів на типовому портативному обладнанні студентського класу. Основна гіпотеза полягала в тому, що композитні моделі забезпечують прийнятну якість прогнозів за критично низького споживання ресурсів, недоступного для сучасних *LLM*-подібних структур.

Таблиця 1. Матриця функційної та технічної спроможності архітектур НМ

Критерій оцінювання	<i>RNN / LSTM</i>	<i>GNN</i>	<i>Transformers</i>
Темпоральна точність	Висока. Оптимально для динамічного трекінгу.	Низька. Потребують темпоральної гібридизації.	Висока. Глобальне моделювання контексту.
Структурна валідність	Низька. Втрата топології через лінеаризацію.	Дуже висока. Апріорні структурні знання.	Середня. Імпліцитне вивчення структури.
Семантична гнучкість	Середня. Інформаційне вузьке місце в текстах.	Низька. Специфічні для завдань класифікації.	Дуже висока. Нелокальні залежності.
Обчислювальна складність	$O(N)$. Лінійна. Ефективні для локальних систем.	$O(\ V\ + \ E\)$. Залежить від щільності.	$O(N^2)$. Квадратична. Вимагають <i>GPU</i> .

Технічні умови та інструментарій. Експериментальний стенд базується на ноутбуці з CPU Intel Core i5-1135G7 (4 ядра, 2,4–4,2 ГГц) та дискретним відеоадаптером NVIDIA GeForce MX350 на архітектурі Pascal. Карта має лише 2 ГБ GDDR5 VRAM, що є критичним порогом для сучасного глибокого навчання. Програмний стек: Ubuntu 22.04, Python 3.10, PyTorch 2.1.0 з підтримкою CUDA 11.8.

Для моделювання реальних умов багатозадачності обсяг доступної відеопам'яті (VRAM) для тренувального та інференс-процесів було програмно обмежено до 1 ГБ. Це імітує ситуацію, коли частина ресурсів зайнята операційною системою та інфокомунікаційними клієнтами.

Теоретичний аналіз FLOPs:

1. Монолітний Трансформер: Модель типу *Tiny-BERT* (4 шари, $d = 312$, де d — розмірність прихованого ембедінгу). За довжини послідовності $N = 512$, обчислення механізму *Self-Attention* потребує приблизно $2 \times N^2 \times d$ операцій (охоплюючи як оцінки QK^T , так і зважену суму V). Тільки для матриць *Attention* (512×512) за 32-бітного формату із плаваючою комою одинарної точності необхідно 1 МБ на кожен шар уваги. З урахуванням активацій та ваг (45 мільйонів параметрів), теоретичний мінімум пам'яті для інференсу одного батча становить $\approx 680\text{--}850$ МБ.

2. Композитна модель: Поеднання блоку *LSTM* (128 одиниць) та 2-шарової *GCN* (*Graph Convolutional Network*). Кількість параметрів — 0,78 мільйона. Теоретична складність операції *LSTM* становить $O(N \times (8d^2 + 8dh))$, де d — розмірність вхідних ознак, а h — кількість прихованих одиниць. При $N = 512$, це потребує у $\approx 18\text{--}20$ разів менше FLOPs, ніж Трансформер.

Експериментальна процедура містила стандартизовану фазу *warm-up* з 50 ітерацій для усунення накладних витрат на ініціалізацію та шум кешування CUDA. Потім кожна модель виконувала безперервний цикл з 1000 циклів інференсу, імітуючи предиктивний потік у реальному часі для цифрового класу. Використовувалися високоточні лічильники продуктивності (*time.perf_counter*) для фіксації затримки інференсу, а бібліотека *NVIDIA Management Library* (NVML) забезпечувала опитування розподілу VRAM, частоти ядра GPU та теплового стану в реальному часі.

Результати вимірювань швидкості та стабільності системи наведено у Таблиці 2.

Аналіз апаратних аномалій:

1. Проблема *Out of Memory*: за умов збільшення довжини історії навчання до 1024 кроків Трансформер миттєво викликав помилку *RuntimeError: CUDA out of memory*, оскільки квадратичне зростання матриць уваги перевищило ліміт 1 ГБ. Композитна модель працювала бездоганно, використовуючи лише 58 МБ пам'яті.

2. Обчислювальний тротлінг: через високу обчислювальну інтенсивність Трансформера, після 4 хвилин безперервного прогнозуван-

ня частота ядра GPU MX350 впала з 1468 МГц до 920 МГц, що призвело до збільшення затримки з 86 мс до 118 мс. Композитна модель не викликала перегріву, зберігаючи пікову частоту протягом усього тесту.

3. Енергоефективність: оціночне енергоспоживання (*TDP*) під час роботи композитної архітектури було на 40% нижчим від номінального. Це критично для роботи ноутбука від акумулятора в умовах відсутності електропостачання, забезпечуючи до 3,5 годин додаткової автономності під час розрахунків.

Висновки розділу. Експеримент довів, що для реальних умов експлуатації на “студентському” обладнанні використання гібридних архітектур є не просто перевагою, а технічною умовою стабільності. Втрата точності у розмірі 0,009 *AUC* (менше 1,1%) є абсолютно виправданою платою за 14-кратне зростання швидкодії та можливість паралельного оброблення сотень потоків даних без ризику перегріву чи зупинки системи через брак пам’яті.

Адаптивні стратегії для оптимізації освітніх траєкторій

Апаратна стабільність та здатність композитних моделей функціонувати на локальних пристроях початкового рівня, перевірені під час експерименту, стають технологічним фундаментом для впровадження адаптивних стратегій у специфічних умовах українського освітнього простору. Воєнні дії та системна нестабільність критичної інфраструктури створили потребу в інструментах, які не залежать від постійного зв’язку з хмарою. Забезпечуючи високу швидкість (6,1 мс на інференс) та енергоефективність, запропоновані архітектури стають основою для підтримки когнітивної стійкості учнів [12, 13]. Очікується, що імплементація таких систем на базі *LSTM* скоротить “освітні розриви” серед внутрішньо переміщених осіб, надаючи можливість навчатися автономно під час блекаутів.

Практична реалізація цих стратегій базується на використанні *LLM* як інтелектуальних асистентів, проте саме композитний рівень

Таблиця 2. Емпіричні дані продуктивності
на ноутбуці (MX350, ліміт 1 ГБ VRAM)

Параметри порівняння	Монолітний Трансформер	Композитна модель	Різниця ефективності
Споживання VRAM (Ваги+Кеш)	914,2 МБ (91,4%)	42,8 МБ (4,3%)	-95,3%
Час інференсу (CPU), мс	248,5	15,2	-93,8%
Час інференсу (GPU), мс	86,4	6,1	-92,9%
Максимальний розмір батча	2 (<i>Out of Memory</i> за 4)	128+	>64x
Температура GPU через 10 хв, °C	78° (Тротлінг)	54° (Стабільно)	-24 °C
Точність (<i>AUC ROC</i>)	0,841	0,832	-0,009

(локальне оброблення) відповідає за адаптивний розподіл підзадач відповідно до актуального стану учня [15]. Процес прийняття рішень математично формалізується через навчання з підкріпленням (*Reinforcement Learning, RL*), де вибір оптимальної педагогічної дії детермінується політикою π (8):

$$\pi(a|s) = P(A_t = a | S_t = s), \quad (8)$$

де $\pi(a|s)$ — стохастична політика агента (нейромережі), що визначає розподіл ймовірностей дій; P — умовна ймовірність настання події; A_t — випадкова величина, що позначає обрану дію (навчальне завдання) у момент часу t ; S_t — випадкова величина, що характеризує поточний когнітивний стан учня; a — конкретна реалізація дії з простору допустимих педагогічних інтервенцій; s — векторний опис поточного рівня компетентності учня.

Стратегічна мета *RL*-системи полягає у максимізації інтегральної винагороди $J(\pi)$ (9):

$$J(\pi) = E_{\pi} \left[\sum_{t=0}^{\infty} \gamma^t r_t \right], \quad (9)$$

де J — цільова функція; E — математичне сподівання; γ — коефіцієнт дисконтування ($0 < \gamma < 1$), що визначає важливість майбутніх винагород; r_t — отримана винагорода (*reward*) на кроці t .

Ефективним технологічним розв'язанням проблеми дефіциту національних наборів даних є впровадження методів крослінгвальних ембедінгів (10):

$$V_{uk} = W \cdot V_{en}, \quad (10)$$

де V_{uk} — векторний простір слів української мови; V_{en} — векторний простір слів англійської мови (багатої на ресурси); W — матриця лінійної трансформації, отримана через методи вирівнювання латентних просторів [16].

Цей підхід дозволяє проєктувати англійські знання на український контент, забезпечуючи лінгвістичну валідність системи навіть за дефіциту локальних наборів даних.

Етика та епістемологія

Технологічна ефективність та перенесення обчислень на пристрої користувачів, підтверджені експериментально, актуалізують не лише питання швидкодії, а й глибинні етико-епістемологічні виклики. Здатність композитних моделей працювати локально на 1 ГБ VRAM створює умови для захисту приватності даних «за дизайном» (*Privacy by Design*), оскільки персональні когнітивні профілі учнів можуть оброблятися без передавання на зовнішні сервери. Водночас, використання гібридного підходу (*LSTM+GNN*) замість монолітних Трансформерів частково вирішує проблему «чорної скриньки», оскільки графові структури є більш інтерпретованими для педагогів.

Принципи прозорості та справедливості, викладені в [17, 18], вимагають, щоб висока швидкість інференсу не досягалася ціною алгоритмічного упередження. Мінімізація ризиків сегрегації учнів реалізується через багатокomпонентну функцію втрат (11):

$$Loss = \lambda_1 L_{accuracy} + \lambda_2 L_{fairness} + \lambda_3 L_{privacy}, \quad (11)$$

де $L_{accuracy}$ — стандартна помилка прогнозування (точність); $L_{fairness}$ — штраф за нерівність у прогнозах для різних демографічних груп; $L_{privacy}$ — компонент, що відповідає за збереження приватності (наприклад, через диференційну приватність); $\lambda_{1,2,3}$ — вагові коефіцієнти, що визначають баланс між ефективністю та етичністю.

Надмірна автоматизація, навіть за високої технічної стабільності, несе загрозу когнітивної гетерономії [19, 20]. Отже, впровадження НМ в Україні має супроводжуватися розробкою етичних регламентів, які гарантують, що ШІ залишається підсилювачем людських можливостей, а не інструментом для деградації критичного мислення.

Висновки

Комплексне наукове обґрунтування та емпірична валідація, проведені в цьому дослідженні, підтверджують, що оптимізація освітніх траєкторій у ресурсообмежених інфокомунікаційних середовищах безпосередньо залежить від архітектурного синтезу нейронних мереж. Порівняльний аналіз встановив, що хоча монолітні моделі Трансформерів забезпечують краще семантичне розуміння, їхня квадратична обчислювальна складність $O(N^2)$ та масивний обсяг пам'яті (заповнення 91,4 % бюджету VRAM об'ємом 1 ГБ лише вагами) роблять їх непридатними для масового розгортання на учнівському обладнанні початкового рівня. Навпаки, синтезована композитна архітектура ($LSTM+GNN$) використовує переваги лінійної складності $O(N)$ рекурентних блоків та топологічну ефективність графових згорток, досягаючи функційного балансу між предиктивною точністю та експлуатаційною стабільністю.

Експериментальні результати надають вагомі докази технічної переваги гібридного підходу в реальних сценаріях. Зафіксоване 14-кратне скорочення затримки інференсу на GPU (з 86,4 мс до 6,1 мс) та 20-кратне зниження споживання відеопам'яті є не просто кількісними покращеннями, а якісними зрушеннями, які дають змогу створювати інтерактивні освітні сервіси в реальному часі на споживчих ноутбуках. Важливо, що композитна модель повністю обійшла помилки *CUDA Out-of-Memory* та термальний тротлінг (78° C проти 54° C), які катастрофічно погіршували продуктивність базової монолітної моделі. Ці висновки свідчать, що для великомасштабних освітніх платформ композитні архітектури є єдиним життєздатним шляхом до забезпечення доступності системи без постійної залежності від централізованої хмарної інфраструктури.

Технічно інтеграція таких моделей безпосередньо відповідає стратегічним потребам українського освітнього сектору. Продемонстроване зниження енергоспоживання на 46 % та здатність ефективно функціонувати на локальних пристроях користувачів створюють основу для когнітивної стійкості під час періодів нестабільності інфраструктури або повної втрати зв'язку. Застосування методів крослінгвальних ембедінгів додатково долає розрив у ресурсах між глобальними англомовними знаннями та локальним українським контентом, гарантуючи, що технічна оптимізація не відбувається за коштом культурної чи лінгвістичної валідності.

З етичної та епістемологічної точки зору, перехід до композитних архітектур сприяє впровадженню принципів «*Privacy by Design*». Завдяки забезпеченню високопродуктивної обробки на 1 ГБ відеопам'яті, запропонована система дає змогу конфіденційним когнітивним профілям залишатися на пристрої користувача, мінімізуючи ризики витоку даних та спостереження. Крім того, притаманна графовим нейронним мережам інтерпретованість пропонує рішення дилеми «чорної скриньки» глибокого навчання, даючи змогу освітянам розуміти структурну логіку алгоритмічних рекомендацій. Підсумовуючи, синтез композитних архітектур є людиноцентричною технологічною еволюцією, яка пріоритезує стабільність, автономність та етичну підзвітність, гарантуючи, що штучний інтелект виступає надійним когнітивним підсилювачем як для учнів, так і для вчителів у цифрову епоху.

ПОВІДОМЛЕННЯ

Внесок авторів: Під час підготовки цієї статті автори використовували велику мовну модель *Gemini 3.1 Pro*. Інструмент застосовувався як допоміжний засіб для структурування матеріалу, стилістичного редагування тексту та допомоги у формулюванні окремих тез. Після використання систем штучного інтелекту автор ретельно перевіряв і відредагував згенерований контент. Автори беруть на себе повну відповідальність за кінцевий зміст і достовірність поданої роботи.

ЛІТЕРАТУРА / REFERENCES

1. Mukul E., Büyüközkan G. Digital transformation in education: A systematic review of education 4.0. *Technological Forecasting and Social Change*, 2023, Vol. 194, Article. 122664. <https://doi.org/10.1016/j.techfore.2023.122664>
2. Vakhabova A.S., Kosulin V.V., Zizaeva A.A. Artificial intelligence in education: Challenges and opportunities for sustainable development. *Economics and Management Problems Solutions*, 2025, Vol. 158, 173-179. <https://doi.org/10.36871/ek.up.p.r.2025.05.09.020>
3. Pinchuk O.P., Malyska I.D. Responsible and ethical use of artificial intelligence in scientific research and publications. *Information Technologies and Learning Tools*, 2024, Vol. 100 (2), 180–198. [In Ukrainian: Пінчук О.П., Малицька І.Д. Відповідальне та етичне використання штучного інтелекту в наукових дослідженнях та публікаціях] <https://doi.org/10.33407/itlt.v100i2.5676>

4. Zawacki-Richter O., Marín V.I., Bond M., Gouverneur F. Systematic review of research on artificial intelligence applications in higher education – where are the educators? *International Journal of Educational Technology in Higher Education*, 2019, Vol. 16, Article 39. <https://doi.org/10.1186/s41239-019-0171-0>
5. Vaswani A. Attention Is All You Need. *Advances in Neural Information Processing Systems (NIPS)*, 2017, Vol. 30, 5998–6008. <https://doi.org/10.48550/arXiv.1706.03762>
6. LeCun Y., Bengio Y., Hinton G. Deep learning. *Nature*, 2015, Vol. 521, 436–444. <https://doi.org/10.1038/nature14539>
7. Makohon Y.O. Use of neural network structures for anomaly detection in cyberspace. *Proc. Intern. Sci. Conf. ICTC-2025, NURE, Kharkiv*, 2025, pp. 306–307. [In Ukrainian: Макогон Ю.О. Використання нейромережових структур для виявлення аномалій у кіберпросторі]
8. Chen X., Xie H., Zou D., Hwang G.J. Application and theory gaps during the rise of Artificial Intelligence in Education. *Computers and Education: Artificial Intelligence*, 2020, Vol. 1, Article 100002. <https://doi.org/10.1016/j.caeai.2020.100002>
9. Scarselli F. The Graph Neural Network Model. *IEEE Transactions on Neural Networks*, 2009, Vol. 20 (1), 61–80. <https://doi.org/10.1109/TNN.2008.2005605>
10. Jobin A., Ienca M. The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 2019, Vol. 1, 389–399. <https://doi.org/10.1038/s42256-019-0088-2>
11. Graves A., Mohamed A., Hinton G. Speech recognition with deep recurrent neural networks. *IEEE International Conference on Acoustics, Speech and Signal Processing*, 2013, 6645–6649. <https://doi.org/10.1109/ICASSP.2013.6638947>
12. Wu Z. A Comprehensive Survey on Graph Neural Networks. *IEEE Transactions on Neural Networks and Learning Systems*, 2021, Vol. 32 (1), 4–24. <https://doi.org/10.1109/TNNLS.2020.2978386>
13. Knox J. Artificial intelligence and education in China. *Learning, Media and Technology*, 2020, Vol. 45 (3), 298–311. <https://doi.org/10.1080/17439884.2020.1754236>
14. Marienko M.V., Kovalenko V.V. Artificial intelligence and open science in education. *Physical and Mathematical Education*, 2023, No. 38(1), 48–53. [In Ukrainian: Марієнко М.В., Коваленко В.В. Штучний інтелект та відкрита наука в освіті] <https://doi.org/10.31110/2413-1571-2023-038-1-007>
15. Crompton H., Bernacki M., Greene J.A. Psychological foundations of emerging technologies for teaching and learning in higher education. *Current Opinion in Psychology*, 2020, Vol. 36, 101–105. <https://doi.org/10.1016/j.copsyc.2020.04.011>
16. Bahdanau D., Cho K., Bengio Y. Neural Machine Translation by Jointly Learning to Align and Translate. *International Conference on Learning Representations (ICLR)*, 2015. <https://doi.org/10.48550/arXiv.1409.0473>
17. Floridi L., Cowls J. A Unified Framework of Five Principles for AI in Society. *Harvard Data Science Review*, 2019, Vol. 1 (1). <https://doi.org/10.1162/99608f92.8cd550d1>
18. Hochreiter S., Schmidhuber J. Long Short-Term Memory. *Neural Computation*, 1997, Vol. 9 (8), 1735–1780. <https://doi.org/10.1162/neco.1997.9.8.1735>
19. Roll I., Wylie R. Evolution and Revolution in Artificial Intelligence in Education. *International Journal of Artificial Intelligence in Education*, 2016, Vol. 26, 582–599. <https://doi.org/10.1007/s40593-016-0110-3>
20. Siemens G. Learning Analytics: The Emergence of a Discipline. *American Behavioral Scientist*, 2013, Vol. 57 (10), 1380–1400. <https://doi.org/10.1177/0002764213498851>

Отримано / Received 05.01.2026

Прийнято / Accepted 24.06.2026

Опубліковано / Published 00.00.2026

Yu.O. MAKOHON, PhD Student,
Kharkiv National University of Radioelectronics,
14, Nauky ave., Kharkiv, 61166, Ukraine
<https://orcid.org/0009-0005-4369-1555>; yurii.makohon@nure.ua

O.H. RUDENKO, DSc (Engineering), Professor,
Kharkiv National University of Radioelectronics,
14, Nauky ave., Kharkiv, 61166, Ukraine
<https://orcid.org/0000-0003-0859-2015>; oleh.rudenko@nure.ua

ARCHITECTURAL SYNTHESIS OF COMPOSITE NEURAL NETWORK SYSTEMS FOR PREDICTIVE DATA PROCESSING IN DIGITAL EDUCATIONAL ENVIRONMENTS

Introduction. The ongoing digital transformation of global educational environments necessitates the deployment of intelligent frameworks capable of real-time optimization of individual learning pathways. Neural networks (NN) have emerged as primary technological drivers in this process, providing advanced predictive analytics of learning activities. However, the increasing parametric complexity of modern monolithic architectures, such as Transformers, creates significant technical barriers for their widespread use on entry-level student hardware (laptops and tablets). In the Ukrainian context, where infrastructure instability often limits access to reliable high-performance cloud resources, the synthesis of lightweight yet highly accurate composite models capable of local data processing is a critical strategic requirement.

The purpose of the paper is to substantiate the methodological and technical principles for synthesizing composite neural network architectures (integrating RNN, GNN, and Transformers) to intensify adaptive learning processes and ensure stable predictive data processing in resource-constrained environments.

Methods. The study employs a formal framework of non-linear dynamic systems to model student cognitive progress. The architectural synthesis involves a modular task decomposition strategy: temporal dynamics are managed by LSTM blocks, structural knowledge dependencies are captured by GCN layers, and semantic interactions are processed via self-attention mechanisms. Experimental validation was conducted on a laptop equipped with consumer-grade hardware (Intel i5-1135G7 CPU, NVIDIA MX350 GPU) with programmatically restricted video memory (1 GB) to simulate multitasking environments. Performance metrics included high-precision latency counters, hardware monitoring via NVML, and predictive accuracy assessment using AUC ROC.

Results. Theoretical FLOPs analysis indicated that the proposed composite architecture requires 18–20 times fewer floating-point operations than monolithic Transformer baselines. For a sequence length of $N = 512$, calculating the Self-Attention mechanism requires approximately $2 \times N^2 \times d$ operations (covering both OK^T and V). In contrast, the LSTM-based composite requires $O(N \times (8dh))$ operations, providing massive efficiency gains. Empirical testing recorded a 14-fold reduction in GPU inference latency (from 86.4 ms down to 6.1 ms) and a 20-fold reduction in the memory footprint. The architecture successfully bypassed CUDA Out-of-Memory errors and prevented thermal throttling, maintaining a stable GPU operating temperature (54° C vs 78° C).

Conclusions. The optimization of educational trajectories is strictly dependent on the architectural synthesis of neural networks. Monolithic models are often impractical for local use due to their $O(N^2)$ complexity and high memory consumption. Composite ensembles provide a necessary balance between predictive precision and operational stability. The ability to process data locally on user devices ensures both educational continuity during infrastructure instability and data privacy by design. These findings establish that hybrid, resource-efficient architectures are the most viable path for creating sustainable, human-centric educational ecosystems in the digital era.

Keywords: *adaptive learning, composite neural network architectures, infocommunication systems, predictive data processing, personalization, local computing, artificial intelligence.*