

Завгородня Г.А.Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»**Завгородній В.В.**Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»

МОДЕЛЮВАННЯ ПОВЕДІНКИ ГРАВЦЯ ЧЕРЕЗ НЕЙРОМЕРЕЖЕВІ АГЕНТИ

У сучасних відеоіграх важливим завданням є створення ігрових агентів, здатних відтворювати індивідуальний стиль гри людини. Це забезпечує більш реалістичну поведінку супротивників, покращує занурення гравця у віртуальне середовище та відкриває можливості для персоналізованого геймплею. У статті представлено підхід до розробки нейромережевого агента, який імітує стиль гри конкретного користувача на основі комбінованого застосування імітаційного та підкріплювального навчання.

Запропонована модель використовує векторне представлення індивідуальних характеристик гравця, що дозволяє формалізувати його стиль через систему параметрів, пов'язаних із частотою дій, реакцією на події та рівнем ризику. Використання алгоритмів кластеризації та нейронних мереж дає змогу відтворювати поведінкові патерни навіть за обмеженої кількості тренувальних даних. Етап підкріплювального навчання забезпечує адаптацію агента до нових ситуацій, дозволяючи йому удосконалювати власну стратегію відповідно до змін у середовищі гри.

Експериментальні дослідження, проведені на основі набору ігрових сценаріїв, показали, що запропонований підхід перевершує традиційні методи імітації за точністю та стабільністю. Зокрема, середнє відхилення поведінкових метрик між агентом і людиною зменшено більш ніж на 15 %, а рівень успішності адаптації до нових умов зріс на 20 % порівняно з базовими моделями. Підхід також продемонстрував ефективність у реальному часі завдяки оптимізації обчислень та використанню спрощених архітектур нейронних мереж.

Отримані результати свідчать, що розроблений нейромережевий агент здатний не лише відтворювати стиль гри людини, а й адаптуватися до змін у її поведінці. Це створює підґрунтя для впровадження персоналізованих систем штучного інтелекту в сучасні ігрові рушії та відкриває нові напрями досліджень у сфері когнітивного моделювання гравців.

Ключові слова: нейромережевий агент, імітація стилю гри, глибинне навчання, підкріплювальне навчання, моделювання поведінки гравця, рекурентні нейромережі, адаптивні ігрові системи.

Постановка проблеми. Однією з ключових проблем сучасного ігрового штучного інтелекту є створення агентів, здатних відтворювати індивідуальний стиль гри конкретного гравця. Традиційні алгоритми на основі правил або класичних методів підкріплювального навчання часто не враховують нюанси поведінки людини та її здатність адаптуватися до динамічних умов гри [1–3]. Відтворення стилю гравця стає критично важливим не лише для підвищення реалізму ігрових симуляторів, а й для розробки тренувальних систем, персоналізованих ігрових сценаріїв та інтелектуальних супротивників [4, 5].

Однією з проблем є моделювання послідовностей дій гравця, що включає як стратегічні

рішення, так і швидкі реакції на події в грі. Стандартні алгоритми часто не здатні точно передбачити наступні кроки користувача, особливо у складних багатокрокових ігрових середовищах [6–8]. Використання нейромережевого підходу дозволяє врахувати історію поведінки, контекст та індивідуальні особливості гравця, що робить імітацію більш природною та реалістичною [9, 10].

Водночас, інтеграція нейромережевих моделей у реальні ігрові рушії пов'язана з обмеженнями продуктивності та обчислювальними ресурсами, що створює додаткові виклики при реалізації реального часу [11, 12]. Таким чином, існує потреба у розробці ефективних моделей нейроме-

режеских агентів, здатних імітувати стиль гравця з високою точністю, враховуючи як поведінкові дані, так і обмеження апаратного забезпечення, а також використання процедурних алгоритмів для генерації контенту та математичного моделювання ігрового середовища [5, 10].

Аналіз останніх досліджень і публікацій. Останні роки спостерігається активний розвиток досліджень, спрямованих на створення нейромережеских агентів, здатних імітувати поведінку людини у відеоіграх. Основна увага приділяється застосуванню методів імітаційного навчання (*imitation learning*), які дозволяють агентам навчатися, відтворюючи дії реальних гравців на основі зібраних траєкторій поведінки [5, 13, 14]. Цей підхід довів ефективність у середовищах, де цільова функція або правила гри не повністю відомі, що робить його придатним для моделювання людського стилю гри.

Сучасні роботи також використовують глибокі рекурентні нейронні мережі (*LSTM*, *GRU*) для збереження контексту попередніх дій і прогнозування наступних рішень гравця [7, 14]. Зокрема, дослідження [9] показало, що поєднання імітаційного навчання з підкріпленням (*reinforcement learning*) дозволяє отримати агентів, які не лише копіюють дії людини, а й адаптуються до нових сценаріїв гри.

Додаткову увагу приділяють побудові моделей персоналізації – алгоритмів, що враховують індивідуальні характеристики гравців, їх стиль прийняття рішень, рівень ризику та реактивність [2, 4, 5]. Методи кластеризації та байєсівського моделювання використовуються для групування гравців за стилем поведінки, що дає змогу створювати більш реалістичних агентів [6, 10]. Крім того, з'являються дослідження з використання генеративних моделей (*GANs*, *diffusion models*) для синтезу нових ігрових дій, які відповідають людській поведінці [3].

Сучасні роботи також демонструють значущість математичного моделювання та формальних методів дослідження для оптимізації агентів та контролю рівня складності [8]. Крім того, використання алгоритмів процедурної генерації контенту дозволяє адаптувати ігрове середовище під конкретного гравця, що підвищує ефективність адаптивної системи [5]. Окрему увагу приділяють застосуванню методів машинного навчання для пошуку аномалій у даних та динамічної корекції поведінки агента у реальному часі [12].

Незважаючи на значний прогрес, проблема створення універсальної архітектури агента, здатного

імітувати стиль гри різних користувачів у режимі реального часу, залишається відкритою. Це зумовлює актуальність подальших досліджень у напрямі поєднання глибокого навчання, процедурного генеративного підходу, поведінкової аналітики та оптимізації обчислювальних ресурсів [5, 10, 12].

Постановка завдання. Метою даного дослідження є розробка нейромережеского агента, здатного імітувати індивідуальний стиль гри людини у відеоігрових середовищах. Результатом дослідження повинна стати інтегрована нейромережеска система, яка здатна відтворювати індивідуальний стиль гри користувача та слугувати основою для персоналізованих ігрових середовищ, інтелектуальних супротивників та тренувальних симуляторів.

Виклад основного матеріалу. Для імітації поведінки гравця ми використовуємо комбіновану нейромережеску модель, що складається з трьох основних компонентів:

1. Мережа прогнозування дій (*Action Prediction Network, APN*) – рекурентна нейромережа (*LSTM*), яка на основі історії дій гравця прогнозує наступну дію a_t .

2. Мережа оцінки стану (*State Evaluation Network, SEN*) – глибока нейромережа, що оцінює стан гри s_t з точки зору стилю та стратегії гравця.

3. Мережа оптимізації дій (*Policy Network, PN*) – підкріплювальна нейромережа, що на основі прогнозів *APN* та оцінок *SEN* визначає оптимальні дії агента.

Структурно агент описується як трійка (S, A, π_θ) , де:

S – множина ігрових станів;

A – множина можливих дій;

$\pi_\theta : S \rightarrow P(A)$ – політика агента з параметрами θ , що задає ймовірність вибору кожної дії.

Нехай $\tau = \{(s_0, a_0), (s_1, a_1), \dots, (s_T, a_T)\}$ – послідовність станів та дій гравця під час сесії. Завдання нейромережеского агента полягає у мінімізації відстані між його діями a_t та реальними діями гравця a_t :

$$\mathcal{L}_{\text{imitation}} = \frac{1}{T} \sum_{t=0}^T l(a_t, a_t),$$

де T – загальна кількість часових кроків, протягом яких порівнюється поведінка гравця і агента;

$l(a_t, a_t)$ – функція втрат (крос-ентропія для дискретних дій).

Для підвищення адаптивності до довгострокових стратегій використовуємо підкріплювальне навчання із винагородою $R(s_t, a_t)$, яка враховує схожість дій агента зі стилем гравця та ефективність у грі:

$$J(\theta) = \mathbb{E}_{\tau \sim \pi_\theta} \left[\sum_{t=0}^T \gamma^t R(s_t, a_t) \right],$$

де $\gamma \in [0, 1]$ – коефіцієнт дисконтованої винагороди.

Алгоритм навчання агента:

1. Збір траєкторій дій гравців τ .
2. Нормалізація станів, категоризація дій, кодування додаткових параметрів (швидкість реакцій, таймінги, параметри стратегії).
3. Імітаційне навчання APN , мінімізація $\mathcal{L}_{imitation}$ методом *Adam*:

$$\theta \leftarrow \theta - \eta \nabla_{\theta} \mathcal{L}_{imitation},$$

де η – швидкість навчання.

4. Підкріплюване навчання PN , оптимізація політики агента через *Policy Gradient*:

$$\nabla_{\theta} J(\theta) = \mathbb{E}_{\tau \sim \pi_{\theta}} \left[\sum_{t=0}^T \nabla_{\theta} \log \pi_{\theta}(a_t | s_t) R(s_t, a_t) \right].$$

5. Інтеграція SEN , оцінка стану $V_{\phi}(s_t)$ через глибоку нейромережу для стабілізації навчання:

$$\mathcal{L}_{value} = \frac{1}{T} \sum_{t=0}^T (R_t - V_{\phi}(s_t))^2.$$

6. Чергування оновлення параметрів θ і ϕ до досягнення заданого порогу точності.

Стиль гравця можна формально описати як вектор характеристик поведінки:

$$\mathbf{p} = [\bar{a}, \sigma_a, \bar{t}, \sigma_t, k_s],$$

де $\bar{a}$ – середнє значення вибору дій;

σ_a – дисперсія дій;

$\bar{t}, \sigma_t$ – середнє та дисперсія часу реакції;

k_s – коефіцієнт ризику або агресивності у стратегії.

Агент навчається відтворювати $\mathbf{p}$ у кожній ігровій сесії, мінімізуючи різницю:

$$\mathcal{L}_{style} = \mathbf{p}_{agent} - \mathbf{p}_{player}.$$

Для перевірки якості імітації використовується:

1. Точність прогнозування дій:

$$Accuracy = \frac{\# \text{правило передбачених дій}}{\text{Загальна кількість дій}}.$$

2. Схожість стилю – косинусна схожість між векторами характеристик $\mathbf{p}$:

$$CosSim(\mathbf{p}_{agent}, \mathbf{p}_{player}) = \frac{\mathbf{p}_{agent} \cdot \mathbf{p}_{player}}{\|\mathbf{p}_{agent}\| \|\mathbf{p}_{player}\|}.$$

3. Залученість користувача – середній час взаємодії з грою під час тестових сесій.

4. Адаптивність агента – здатність підлаштуватися до змін гравця в режимі реального часу.

Для перевірки ефективності запропонованого нейромережевого агента, що імітує стиль гри людини, було проведено експериментальне порівняння з двома базовими методами:

1. Метод А (*BC – Behavioral Cloning*) – класичне імітаційне навчання без зворотного зв'язку.

2. Метод В (*DQN – Deep Q-Network*) – підкріплювальне навчання без використання поведінкових даних.

3. Метод С (*Proposed Hybrid Model*) – комбінований підхід, який поєднує поведінкове навчання та підкріплювальну оптимізацію дій, використовуючи векторний опис стилю гравця.

Дослідження проводилось у симуляційному середовищі *OpenAI Gym (CarRacing-v2)*, яке дозволяє фіксувати послідовність дій гравців у безперервному просторі станів. Було зібрано 25 000 епізодів поведінки п'яти різних користувачів із різними стилями гри (агресивний, обережний, стратегічний, ризиковий, збалансований).

Для оцінки схожості між агентом та людиною використовувалися такі метрики:

– *Cosine Similarity (CS)* між вектором дій агента і гравця;

– *Fréchet Distance (FD)* між послідовностями траєкторій;

– *Success Rate (SR)* – частка завершених епізодів без аварій;

– *Adaptation Time (AT)* – середня кількість кроків, необхідних агенту для пристосування до нового середовища.

Результати експериментів наведені в таблиці 1.

Отримані результати свідчать, що запропонований підхід перевершує традиційні методи за всіма ключовими показниками. Поєднання імітаційного навчання (для збереження стилю гравця) та підкріплювального навчання (для оптимізації дій у нових умовах) забезпечило:

– зростання схожості поведінки з реальними гравцями на 16–25% у порівнянні з базовими методами;

Таблиця 1

Порівняльні результати роботи нейромережевого агента та базових моделей

№	Метод	CS	FD	SR (%)	AT (кроків)	Примітка
1	Behavioral Cloning (BC)	0.78	1.32	72.4	1800	швидко навчається, але не адаптується
2	Deep Q-Network (DQN)	0.65	1.85	68.1	2500	потребує великої кількості епізодів
3	Proposed Hybrid Model (BC+RL)	0.91	0.94	88.7	950	найкраща схожість та адаптація

– зниження часу адаптації на понад 40%, що критично для інтерактивних ігор у реальному часі;
 – підвищення стабільності при переході між різними рівнями складності гри, завдяки функції втрат, що враховує часову узгодженість стилю.

Ключовим компонентом моделі став векторний опис стилю гравця, що дозволяє класифікувати агентів за типом поведінки та використовувати його як вхід для трансформерної моделі, що прогнозує наступну дію. Крім того, запропонована архітектура може бути інтегрована в *Unity*, *Unreal Engine* або *Godot* завдяки використанню оптимізованої інференс-моделі на основі *ONNX Runtime*, що забезпечує виконання в реальному часі без значного навантаження на процесор чи *GPU*.

Висновки. У результаті проведеного дослідження було розроблено та експериментально перевірено гібридну модель нейромережевого агента, здатного імітувати індивідуальний стиль гри людини. Запропонований підхід базується на інтеграції двох парадигм машинного навчання – імітаційного та підкріплювального, що дозволяє поєднати переваги обох: точність відтворення поведінки користувача та здатність до самонавчання в нових ситуаціях.

Створений агент демонструє високу точність поведінкової відповідності (*Cosine Similarity* = 0.91) та знижену фреше-дистанцію між траєкторіями дій гравця та моделі (*FD* = 0.94), що свідчить про реалістичність і стабільність стилю. Крім того, середній показник успішності епізодів (88,7%) перевищує базові методи, а час адаптації скорочується майже вдвічі.

Запропонована векторна репрезентація стилю забезпечує можливість класифікації гравців за поведінковими характеристиками, що відкриває шлях до персоналізації ігрового досвіду. Модель може бути інтегрована у популярні рушії (*Unity*, *Unreal*, *Godot*) без значного навантаження на ресурси, що робить її практично придатною для використання у комерційних і навчальних ігрових системах.

Подальші дослідження планується спрямувати на розширення архітектури агента для мультиплеєрних сценаріїв, удосконалення механізмів самоадаптації та створення мультимодальних моделей стилю, що враховуватимуть не лише дії, а й емоційні та комунікативні аспекти взаємодії користувача з грою.

Список літератури:

1. Chrysafiadi K., Kamitsios M., Virvou M. Fuzzy-based dynamic difficulty adjustment of an educational 3D-game. *Multimedia Tools and Applications*. 2023. Vol. 82. P. 27525–27549. DOI: <https://doi.org/10.1007/s11042-023-14515-w>
2. Romero-Méndez E. A., Santana-Mancilla P. C., García-Ruiz M., Montesinos-López O. A., Anido-Rifón L. E. The Use of Deep Learning to Improve Player Engagement in a Video Game through a Dynamic Difficulty Adjustment Based on Skills Classification. *Applied Sciences*. 2023. 13(14). 8249. DOI: <https://doi.org/10.3390/app13148249>
3. Vang C. The Impact of Dynamic Difficulty Adjustment on Player Experience in Video Games. *Scholarly Horizons: University of Minnesota Morris Undergraduate Journal*. 2022. Vol. 9, No. 1. DOI: <https://doi.org/10.61366/2576-2176.1105>
4. Yannakakis G., Togelius J. Artificial Intelligence and Games. *Cham: Springer*, 2018. 337 p. DOI: <https://doi.org/10.1007/978-3-319-63519-4>
5. Завгородній В.В., Завгородня Г.А., Валявська Н.О., Адаменко В.С., Дороговцев Є.В., Несмачний П. В. Метод автоматичної генерації контенту на основі процедурних алгоритмів. *Вчені записки Таврійського національного університету імені В.І. Вернадського. Серія: Технічні науки*. 2022. Том 33 (72), №1. С. 91–96. DOI: <https://doi.org/10.32838/2663-5941/2022.1/15>
6. Pfau J., Seif El-Nasr M. On Video Game Balancing: Joining Player- and Data-Driven Analytics. *Games Research and Practice*. 2024. Vol. 1, No. 1. P. 1–21. DOI: <https://doi.org/10.1145/3675807>
7. Zheng T. Dynamic difficulty adjustment using deep reinforcement learning: A review. *Applied and Computational Engineering*. 2024. Vol. 71. P. 157–162. DOI: <https://doi.org/10.54254/2755-2721/71/20241633>
8. Завгородній В.В., Завгородня Г.А., Дроботович К.Є., Тенігін О.В., Шматко М.М. Математичне моделювання у методах формального дослідження. *Вчені записки Таврійського національного університету імені В.І. Вернадського. Серія: Технічні науки*. 2021. Том 32 (71), №6. С. 75–79. DOI: <https://doi.org/10.32838/2663-5941/2021.6/12>
9. Rahimi M., Moradi H., Vahabie A.-h., Kebriaei H. Continuous reinforcement learning-based dynamic difficulty adjustment in a visual working memory game. *arXiv preprint*. 2023. DOI: <https://doi.org/10.48550/arXiv.2308.12726>
10. Завгородній В.В., Завгородня Г.А., Демченко І.В., Крамаренко К.С., Шевченко І.О., Юрченко А.В. Метод створення штучних текстур із заданими параметрами. *Вчені записки Таврійського національного*

університету імені В.І. Вернадського. Серія: Технічні науки. 2022. Том 33 (72), №2. С. 86–90. DOI: <https://doi.org/10.32838/2663-5941/2022.2/14>

11. Sepúlveda G. K., Besoain F., Barriga N. A. Exploring Dynamic Difficulty Adjustment in Videogames. *arXiv*. 2020. DOI: <https://doi.org/10.1109/CHILECON47746.2019.8988068>

12. Завгородній В.В., Завгородня Г.А., Валявська Н.О., Герасименко О.О., Калюжний О.В., Степовий А.В. Пошук аномалій у даних за допомогою машинного навчання. *Вчені записки Таврійського національного університету імені В.І. Вернадського. Серія: Технічні науки*. 2022. Том 33 (72), №3. С. 39–43. URL: https://tech.vernadskyjournals.in.ua/journals/2022/3_2022/6.pdf

13. Rodríguez G. R., Neumann U., Heintz F. Player modeling for dynamic difficulty adjustment in top-down shooter games. *Northeastern University Repository*. 2019. URL: https://repository.library.northeastern.edu/files/neu%3Am0455c22w/fulltext.pdf?utm_source=chatgpt.com

14. Mi Q., Gao T. Engagement-oriented dynamic difficulty adjustment (EDDA). *Applied Sciences*. 2025. Vol. 15(10). 5610. DOI: <https://doi.org/10.3390/app15105610>

Zavhorodnia G.A., Zavhorodnii V.V. MODELING PLAYER BEHAVIOR THROUGH NEURAL NETWORK AGENTS

In modern video games, one of the key challenges is developing game agents capable of reproducing the individual playing style of a human player. This ensures more realistic opponent behavior, enhances player immersion in the virtual environment, and enables personalized gameplay experiences. This paper presents an approach to developing a neural network agent that imitates a player's unique style based on a combined application of imitation learning and reinforcement learning.

The proposed model employs a vector representation of individual player characteristics, allowing formalization of the player's style through a system of parameters related to action frequency, response to events, and risk-taking behavior. The use of clustering algorithms and neural networks enables the reproduction of behavioral patterns even with a limited amount of training data. The reinforcement learning stage ensures the agent's adaptation to new situations, allowing it to improve its strategy according to changes in the game environment.

Experimental studies conducted on a set of game scenarios demonstrated that the proposed approach outperforms traditional imitation methods in both accuracy and stability. In particular, the average deviation of behavioral metrics between the agent and human players was reduced by more than 15%, while the adaptation success rate to new conditions increased by 20% compared to baseline models. The approach also proved efficient in real-time applications due to computational optimization and the use of simplified neural network architectures.

The obtained results show that the developed neural network agent is capable not only of reproducing human play style but also of adapting to behavioral changes over time. This provides a foundation for implementing personalized artificial intelligence systems in modern game engines and opens new directions for research in the field of cognitive player modeling.

Key words: neural network agent, gameplay style imitation, deep learning, reinforcement learning, player behavior modeling, recurrent neural networks, adaptive game systems.

Дата надходження статті: 30.10.2025

Дата прийняття статті: 19.11.2025

Опубліковано: 30.12.2025