

Крихівський М. В., кандидат технічних наук, доцент,
доцент кафедри інженерії програмного забезпечення
Івано-Франківського національного технічного університету
нафти і газу
ORCID: 0009-0000-3285-4308

Ваврик Т. О., асистент кафедри інженерії
програмного забезпечення
Івано-Франківського національного технічного університету
нафти і газу, м. Івано-Франківськ
ORCID: 0000-0002-0612-0084

Гобир Л. М., асистент кафедри інженерії
програмного забезпечення,
Івано-Франківського національного технічного університету
нафти і газу
ORCID: 0009-0007-3176-2314

ОСНОВИ ПРОЕКТУВАННЯ ІНТЕЛЕКТУАЛЬНИХ АГЕНТІВ: ВІД АРХІТЕКТУРИ ДО МАТЕМАТИЧНИХ МОДЕЛЕЙ

Інтелектуальні агенти стали центральною темою сучасних досліджень у сфері штучного інтелекту. Вони є важливим компонентом сучасних інформаційних технологій, що використовуються для автоматизації складних завдань, таких як обробка даних, прийняття рішень та адаптивне навчання. Розробка інтелектуальних агентів є складним, багатограним процесом, що включає вибір архітектури, розробку алгоритмів прийняття рішень і інтеграцію з середовищем. Протягом останніх десятиліть з'явилося багато нових технологій та підходів, які дозволяють значно покращити ефективність агентів у реальних застосуваннях. Стаття досліджує концептуальні засади інтелектуальних агентів, різні підходи до їх проектування та ключові технології, що дозволяють створювати автономні та адаптивні системи. Авторами розглядається загальний алгоритм функціонування інтелектуального агента, його основні етапи та архітектура. Представлено огляд концепції інтелектуальних агентів, включаючи їх базові властивості, моделі взаємодії з навколишнім середовищем та адаптивні механізми. Основну увагу приділено універсальному алгоритму, який може бути адаптований до різних застосувань у сфері штучного інтелекту. У статті представлено математичну основу для опису інтелектуальних агентів, що дозволяє моделювати їх поведінку, прийняття рішень та взаємодію з середовищем. У статті представлено математичну основу для опису інтелектуальних агентів, що дозволяє моделювати їх поведінку, прийняття рішень та взаємодію з середовищем. Розглянуто основні компоненти агентів, їх формалізацію через теоретичні моделі та відповідні алгоритми. Також розглянуто математичний опис агентів з урахуванням їхніх властивостей, таких як адаптивність, автономність та взаємодія з середовищем. Математичний аналіз інтелектуальних агентів є важливою і швидко розвиваючою областю, що охоплює широкий спектр методів та технік для моделювання, аналізу та оптимізації поведінки інтелектуальних агентів.

Ключові слова: інтелектуальні агенти, архітектура агентів, класифікація агентів, сенсори, актори.

Krykhivskiy M. V., Vavryk T. O., Hoby L. M. Foundations of Intelligent Agent Design: From Architecture to Mathematical Modeling

Intelligent agents (IA) have become a central topic of modern research in the field of artificial intelligence. They are an important component of modern information technologies used to automate complex tasks such as data processing, decision-making, and adaptive learning. The development of intelligent agents is a complex, multifaceted process that includes the choice of architecture, the development of decision-making algorithms, and integration with the environment. Over the past decades, many new technologies and approaches have emerged that can significantly improve the effectiveness of agents in real-world applications. The article explores the conceptual foundations of intelligent agents, various approaches to their design, and key technologies that allow for the creation of autonomous and adaptive systems. The authors consider the general algorithm of the functioning of an intelligent agent, its main stages and architecture. An overview of the concept of intelligent agents is presented, including their basic properties, models of interaction with the environment and adaptive mechanisms. The main attention is paid

© М. В. Крихівський, Т. О. Ваврик, Л. М. Гобир, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

to a universal algorithm that can be adapted to various applications in the field of artificial intelligence (AI). The article presents a mathematical basis for describing intelligent agents, which allows modeling their behavior, decision-making and interaction with the environment. The article presents a mathematical basis for describing intelligent agents, which allows modeling their behavior, decision-making and interaction with the environment. The main components of agents, their formalization through theoretical models and corresponding algorithms are considered. The mathematical description of agents is also considered, taking into account their properties, such as adaptability, autonomy and interaction with the environment. Mathematical analysis of intelligent agents is an important and rapidly developing area, covering a wide range of methods and techniques for modeling, analyzing and optimizing the behavior of intelligent agents.

Key words: intelligent agents, agent architecture, agent classification, sensors, actors.

Постановка проблеми. ІА відіграють ключову роль у багатьох галузях, включаючи робототехніку, автоматизацію, системи підтримки прийняття рішень та навіть побутові застосування. Розуміння основ проєктування таких агентів допомагає дослідникам та розробникам створювати ефективні та надійні системи, що можуть виконувати складні завдання з мінімальним втручанням людини. Окрім того, стаття також надає важливу теоретичну базу для подальших досліджень і вдосконалення алгоритмів штучного інтелекту (ШІ), що робить її незамінним джерелом знань для науковців і фахівців у цій сфері.

Невирішені питання. Проєктування ІА залишається складним завданням через низку невирішених питань. Однією з головних проблем є забезпечення агента здатністю до розуміння та інтерпретації контексту, адже обробка мови вимагає тонкого розуміння контекстуальних нюансів. Ще одним важливим аспектом є навчання агентів з урахуванням етичних міркувань та уникнення упередженості, що виникає через перекоси в навчальних даних. Крім того, адаптивність агентів до нових, непередбачуваних ситуацій і їх здатність до самонавчання та покращення в реальному часі також викликають занепокоєння. Безпека та конфіденційність даних користувачів залишаються критичними питаннями, які потрібно розв'язувати. Важливим є також забезпечення природної та ефективної взаємодії з користувачами, що включає виявлення та відповідь на емоції. Усі ці питання потребують інноваційних підходів та подальших досліджень для досягнення значного прогресу.

Створення універсальної архітектури інтелектуальних агентів, що поєднує адаптивність, масштабованість і здатність інтегрувати різні підходи, залишається важливим і складним завданням. Однією з ключових проблем є забезпечення агентів здатністю до гнучкого реагування на змінні середовища, що вимагає ефективних механізмів самонавчання та адаптації до нових умов. Це включає в себе використання методів машинного навчання та штучного інтелекту для постійного покращення та коригування поведінки агентів.

Масштабованість є важливим аспектом, оскільки для роботи з великими системами потрібно забезпечити ефективний розподіл ресурсів і обробку даних. Це вимагає розробки архітектури, здатної до горизонтального масштабування, що дозволяє збільшувати потужність системи шляхом додавання нових вузлів або компонентів без зниження продуктивності.

Інтеграція різних підходів, таких як реактивні, розподілені та гібридні, є важливою для створення універсальної архітектури, здатної вирішувати різноманітні завдання. Це потребує забезпечення узгодженості та взаємодії між компонентами системи, що використовують різні підходи, без втрати ефективності. Важливо також враховувати питання безпеки, конфіденційності та етичних міркувань при розробці таких систем. Усі ці аспекти вимагають інноваційних підходів та ретельного планування для досягнення гармонійної та ефективної взаємодії компонентів інтелектуальних агентів.

Пропозиція для вирішення проблеми:

1. Створення окремих модулів для сенсорного сприйняття, прийняття рішень, навчання та взаємодії між агентами.

2. Забезпечення інтерфейсів для легкої інтеграції нових модулів або алгоритмів

3. Запровадження механізмів безперервного навчання: агент оновлює свою модель поведінки в реальному часі, використовуючи досвід.

Очікувані результати:

1. Гнучка та масштабована архітектура, яка адаптується до різноманітних середовищ.

2. Зменшення витрат на інтеграцію нових задач та алгоритмів.

3. Підвищення ефективності агентів у реальному часі.

Аналіз останніх досліджень та публікацій. Методологія BDI (Belief-Desire-Intention), або переконання-бажання-наміри (ПБН), є одним із підходів до моделювання агентів у галузі штучного інтелекту. Вона базується на психологічних теоріях, що пояснюють, як люди приймають рішення, і переносить ці принципи на штучних агентів, що дозволяє їм діяти раціонально та автономно.

Основні компоненти ПБН моделі [1] – це переконання (Beliefs), бажання (Desires) та наміри (Intentions). Переконання – це знання або переконання агента про світ, які можуть бути правдивими або хибними. Це інформація, яка допомагає агенту оцінювати навколишню ситуацію та приймати рішення. Бажання – це цілі або завдання агента, що він хоче досягти. Вони спрямовують його дії та визначають, що агент намагається здійснити. Наміри – це конкретні плани або дії, які агент обирає для досягнення своїх цілей. Це рішення, які агент приймає, виходячи з своїх переконань та бажань.

ПБН архітектура дозволяє агенту працювати у складних динамічних середовищах, здійснюючи розумні та гнучкі рішення. Агент постійно оновлює свої переконання на основі нової інформації, коригує свої бажання, коли змінюються його пріоритети або умови, та обирає нові інтенції, коли це потрібно для досягнення цілей. Це надає агенту здатність адаптуватися до змін та непередбачуваних ситуацій.

Важливою характеристикою ПБН агентів є те, що вони здатні вирішувати конфлікти між різними бажаннями та інтенціями, враховуючи пріоритети та ресурси. Наприклад, якщо агент має кілька суперечливих цілей, він зможе оцінити, яка з них важливіша або більш досяжна у поточний момент, та вибрати відповідну дію. Проте, дослідники [2, 3] відмічають відсутність деяких компонентів гнучкості в моделі BDI та проблемність навчання.

Авторами оглядової статті [3] розглянуті основні підходи до кожного компонента архітектури ПБН, як вони були реалізовані в мовах програмування агентів, і показані компроміси, що властиві кожному підходу. Зроблені висновки про те, що для нових областей, які повинні інтегрувати субсимвольну інформацію від датчиків у переконання агентів, знадобляться більш насичені уявлення, ніж в роботі [1]. Це може передбачати використання вкладених та невизначених переконань.

Довгострокова автономія інтелектуальних агентів (IA) стає необхідною для багатьох автономних систем. Це стимулює подальші дослідження підходів до цілей у контексті того, що цілі підтримання та їх взаємодія з цілями досягнення стають більш важливими, ніж властивості та обмеження цілей, такі як пріоритети та дедлайни.

У роботі [3] виділяється проблема справжнього одночасного виконання планів. Цей напрям передбачає дослідження з формалізації взаємодій між одночасно виконуваними гілками та між циклом міркувань IA та тривалими діями. Інша область майбутніх досліджень стосується використання технік штучного інтелекту, таких як машинне навчання, для доповнення заздалегідь визначених правил плану або генерації нових правил. Авторами підкреслено перспективу підтримки невизначених переконань, що важливо для того, коли декларативна ціль може вважатися «досягнутою» (або недосяжною), і для підходів до планування намірів на основі симуляцій чергувань.

У статті [4] запропоновано архітектуру виконання плану, яка підтримує різну семантику завдань. Кожна мета, що визначена під час циклу обговорення агента, може бути задоволена через завдання, визначені в різній семантиці. Можливість підтримувати різну семантику завдань забезпечує дві основні переваги. Першою перевагою є повторне використання застарілих артефактів у планах агента. Другою перевагою є спрощення адаптації архітектури агента до різних стандартів організації бізнесу. Щоб інтегрувати різноманітну семантику завдання в архітектуру виконання плану, автори [4] використали плавну переглянута версію попередньої сформульованої моделі робочого процесу, у якій різна семантика завдання була б зменшена перед виконанням. Інтегрування ієрархічної мережі завдань і семантики OWL-S в цю архітектуру виконання плану, дозволило перевірити її міцність з точки зору підтримки виконання різних семантик завдань в архітектурі агента.

Авторами роботи [5] створено структуру, яка додає знання для покращення вибору плану в популярній парадигмі програмування агентів ПБН. На відміну від попередніх пропозицій, наведений підхід може добре масштабуватись із складністю бібліотеки планів агента. Технічно розроблено нову міру впевненості, яка дозволяє агенту динамічно регулювати свою залежність від навчання, полегшуючи в принципі нескінченну кількість фаз (повторного) навчання.

Традиційний агент ПБН має 3 основні обчислювальні компоненти, які генерують переконання, генерують наміри та виконують наміри. Вони виконуються послідовно і циклічно. Це може спричинити кілька проблем. Серед них нездатність постійно спостерігати за навколишнім середовищем у динамічних середовищах може бути катастрофічною. Авторами роботи [6] запропоновано загальну структуру для моделі паралельного агента ПБН. За цією загальною структурою можна створювати паралельні агенти ПБН з різними конфігураціями залежно від доступності фізичних ресурсів. Ці агенти мають низку переваг перед послідовними, а саме зміни в середовищі агента можуть бути виявлені негайно, надзвичайні ситуації будуть вирішені негайно, підтримка надається на рівні архітектури для перегляду бажань/намірів і розгляду зв'язків між цілями, коли генерується нове переконання/бажання.

Концепції IA представлені в дослідницькій книзі [7], у якій сформована теорія агентів та адаптацію агентів у різних контекстах. Агенти різного порядку складності повинні бути автономними у використовуваних правилах. Агент повинен мати мозок, за допомогою якого він може виявити правила, що містяться в даних. Оскільки правила є інструментами, за допомогою яких агенти змінюють середовище, будь-яку адаптацію правил можна розглядати як еволюцію агентів. Оскільки невизначеність присутня в будь-якому контексті, авторами книги [7] описано як можна ввести глобальну невизначеність з локального світу в опис правил. У книзі представлено загальний опис агента еволюційної адаптації, дії та мета дії такого агента, абстрактне та формальне представлення дії, зв'язок системи та метасистеми з адаптивним агентом, мозок агента за допомогою теорії морфогенетичних нейронів, логічну структуру адаптивного агента, зворотний зв'язок і гіперзворотний зв'язок в адаптивному агенті, поле адаптації в модальний логічний простір як логічний інструмент в адаптивному агенті, дію агента у фізичному домені та застосування агентів у роботах та еволюційних обчисленнях.

Автори роботи [8] представляють багатоагентну архітектурну теорію для розробки інтелектуальних асоціативних гібридних систем. Теорія архітектури описана як на рівні структури завдання, так і на рівні обчислень. Ця архітектура призначена для вирішення проблем розробки агентів знань та інформаційних агентів.

Архітектури автономних агентів – це методології проектування сукупності знань і стратегій, які застосовуються до проблеми створення локального інтелекту. У статті [9] зроблено спробу об'єднати ці знання з кількома архітектурними традиціями, приділяючи особливу увагу особливостям, які, як правило, вибиралися під тиском широкого використання в системах реального світу. Автори визначили, що наступні стратегії надають значну допомогу в проектуванні автономних інтелектуальних агентів, а саме модульність, яка спрощує як дизайн, так і контроль, ієрархічно організований вибір дій, який фокусує увагу та забезпечує пріоритетність, коли різні модулі конфліктують; та паралельний моніторинг середовища, який дозволяє системі бути чутливою та опортуністичною, дозволяючи переключити увагу та переоцінити пріоритети. Ми пропонуємо огляд чотирьох архітектурних парадигм: III на основі поведінки; дво- і тришарові системи; архітектури віри, бажання та наміру (зокрема PRS); і Soar/ACT-R. Задokumentувавши тенденції в кожній із цих спільнот щодо створення компонентів вище, ми стверджуємо, що ця конвергентна еволюція є переконливим доказом корисності компонентів. Потім ми використовуємо цю інформацію, щоб рекомендувати конкретні стратегії для дослідників, які працюють за кожною парадигмою, для подальшого використання знань і досвіду в галузі в цілому.

Механізми емоцій часто використовуються в штучних агентах як метод покращення вибору дій. Порівняння між агентами є складними через відсутність єдності між теоріями емоцій, завданнями агентів і типами вибору дій, що використовуються. Набір архітектурних якостей пропонується у роботі [10] як основа для проведення порівнянь між агентами. Аналіз існуючих архітектур агентів, які включають механізм емоцій, може допомогти триангулювати можливості дизайну в просторі, окресленому цими якостями. З огляду на це, для аналізу обрано дванадцять автономних агентів, які включають механізм емоцій у вибір дій. Кожен агент розбирається за допомогою цих архітектурних якостей (архітектура агента, механізм вибору дій, механізм емоцій і представлення емоційного стану, а також модель емоцій, на якій він базується). Це допомагає розмістити агентів в архітектурному просторі, підкреслює контрастні методи реалізації схожих теоретичних компонентів і підказує, які архітектурні аспекти важливі для виконання завдань. Представлено початкову структуру, що складається з ряду рекомендацій щодо розробки механізмів емоцій у штучних агентах, заснованих на кореляції між виконуваними ролями емоцій та аспектами механізмів емоцій, які використовуються для виконання цих ролей. У висновку обговорюється, як можна вирішити проблеми з цим типом дослідження та якою мірою розробка основи може допомогти майбутнім дослідженням.

Логіка ПБН є важливим і широко використовуваним теоретичним апаратом для представлення та міркування про раціональних агентів. Однак вона є неповною щодо перегляду наміру, перевизначення наміру, процесу обговорення та перегляду переконань. Крім того, деякі раціональні агенти, особливо люди, не мають наближених міркувань, добре представлених логікою ПБН. У статті [11] визначено нечітку семантику для мови ПБН, яка здатна усунути ці обмеження. Запропоновано розширити логіку ПБН до нечіткої логіки ПБН.

У статтях [12-14] проведено аналіз педагогічних та інформаційних засад використання інтелектуальних агентів в адаптивних системах електронного навчання з метою створення структури адаптивної системи електронного навчання. Авторами описано характеристики, функції та взаємодії агентів, які беруть участь у кожному модулі адаптивної архітектури, а також характеристики, функції та взаємодії інтелектуального агента для ухвалення навчальних рішень.

У роботі [15] представлено підхід до розробки ІА, який можна використовувати в мультиагентній системі для вирішення складної проблеми. Запропонований агент має можливості навчання за допомогою штучних нейронних мереж.

Авторами статті [16] розглядаються проблеми створення надійних децентралізованих програм способом, сумісним із наскрізною моделлю. Створена ними програма Mandrake базується на розумінні того, що можливість агента отримувати повідомлення в будь-якому порядку та обробляти їх відповідно до його власного прийняття рішення є вирішальною для реалізації сенсу програми.

Автори публікації [17] зазначають, що III ніщо без інтелектуальних агентів і досліджують скільки типів агентів визначено в штучному інтелекті. Вони виділяють п'ять категорій, які згруповані за діапазоном можливостей і ступенем сприйманого інтелекту, і роблять висновок, що кожна з них здатна втілити ідеї III в життя.

У дослідженнях [18-21] розшифровано багатогранну природу інтелектуальних агентів, проаналізовано їх застосування в реальному світі, переваги та потенційні обмеження. Автори розглянули класифікацію інтелектуальних агентів, що включає різні типи, такі як прості рефлексивні агенти, модельно-рефлексивні агенти, цільові агенти та корисні агенти. Кожен варіант наділяє агента різними можливостями та функціями, що відповідають різноманітним вимогам додатків III.

Мета статті є надання всебічного огляду та аналізу ключових аспектів створення ІА. Стаття спрямована на забезпечення теоретичної та практичної основи для розуміння процесу розробки таких агентів, починаючи від їхньої архітектури та закінчуючи математичними моделями, які описують їхню поведінку та

функціонування. Це прагнення узагальнити існуючі знання та досвід, а також пропонування нових підходів та методів, що сприятимуть підвищенню ефективності та надійності інтелектуальних систем у різних сферах застосування.

У цій статті розглянуто основні принципи та підходи до проектування ІА, з особливим акцентом на їхню архітектуру та математичне моделювання. Ми прагнемо створити цілісне уявлення про теоретичні основи та практичні аспекти цього процесу, що дозволить розробникам і дослідникам ефективно використовувати ці знання для створення надійних та адаптивних систем.

Виклад основного матеріалу дослідження. Сучасні підходи до проектування інтелектуальних агентів включають низку методологій та технологій, що дозволяють створювати ефективні та гнучкі системи для розв'язання різноманітних завдань. Одним з таких підходів є агентно-орієнтоване програмування, яке базується на концепції автономних агентів, що взаємодіють з оточенням і іншими агентами для досягнення своїх цілей.

У основі сучасного проектування інтелектуальних агентів лежить використання машинного навчання та штучного інтелекту. Одним з ключових аспектів є навчання агентів на основі підкріплення (reinforcement learning), що дозволяє агентам навчатися на власному досвіді через взаємодію з середовищем. Такі агенти використовують алгоритми, що дозволяють їм удосконалювати свої стратегії та приймати оптимальні рішення на основі отриманих винагород.

Ще одним важливим підходом є використання нейронних мереж, зокрема, глибокого навчання (deep learning). Глибокі нейронні мережі дозволяють агентам аналізувати великі обсяги даних і виявляти приховані закономірності, що підвищує їхню здатність до передбачення та прийняття рішень. Такі мережі можуть бути використані для обробки візуальної інформації, розпізнавання мови, а також для розв'язання складних задач, таких як ігри чи стратегічне планування.

Сучасні підходи також передбачають використання мультиагентних систем, де кілька агентів взаємодіють між собою для досягнення спільних цілей. У таких системах агенти можуть співпрацювати, координувати свої дії та обмінюватися інформацією, що дозволяє підвищити ефективність і надійність системи в цілому. Для цього використовуються методи розподіленого штучного інтелекту та алгоритми консенсусу.

Крім того, важливим аспектом проектування інтелектуальних агентів є забезпечення їхньої адаптивності та масштабованості. Сучасні агенти повинні бути здатні адаптуватися до змін у середовищі, а також масштабуватися для роботи в різних умовах і з різним рівнем завантаженості. Для цього використовуються методи самоорганізації та автономного управління.

Отже, сучасні підходи до проектування інтелектуальних агентів базуються на використанні машинного навчання, нейронних мереж, мультиагентних систем і методів адаптивного управління. Ці підходи дозволяють створювати гнучкі, ефективні та надійні системи, здатні розв'язувати складні задачі в різноманітних галузях.

Існує кілька інших методологій для моделювання агентів, кожна з яких має свої особливості та відрізняється від ПБН підходу:

1. Реактивні агенти: Ці агенти працюють на основі безпосередніх реакцій на зміни в навколишньому середовищі. Вони не мають внутрішнього уявлення про світ чи довгострокових цілей. Замість цього, вони реагують на вхідні стимули за допомогою заздалегідь визначених правил або алгоритмів. Це робить їх швидкими та ефективними у простих і стабільних середовищах, але вони можуть бути менш гнучкими у складних та динамічних умовах.

2. Цільові агенти (Goal-oriented agents): Ці агенти працюють на основі цілей, які вони намагаються досягти. Вони використовують планування та розв'язання задач для досягнення своїх цілей. Цей підхід дозволяє агентам бути більш раціональними та адаптивними, оскільки вони можуть планувати свої дії на основі поточної ситуації та ресурсів.

3. Агенти на основі користувацьких моделей (User model-based agents): Ці агенти використовують моделі користувачів для розуміння та прогнозування їх поведінки. Вони можуть адаптувати свої дії та відповіді на основі індивідуальних характеристик та потреб користувачів. Це корисно у застосуваннях, де важлива взаємодія з користувачем, наприклад, у рекомендаційних системах або персональних асистентах.

4. Соціальні агенти: Ці агенти враховують соціальні взаємодії та відносини з іншими агентами. Вони можуть співпрацювати, конкурувати або вести переговори з іншими агентами для досягнення своїх цілей. Соціальні агенти зазвичай використовують моделі багатогентних систем та ігрову теорію для ефективної взаємодії.

5. Гібридні агенти: Цей підхід поєднує елементи кількох методологій для створення більш потужних і гнучких агентів. Наприклад, агент може використовувати реактивні стратегії для швидкої реакції на зміни в середовищі та планування на основі цілей для довгострокових завдань. Гібридні агенти можуть мати складну архітектуру, але вони здатні краще адаптуватися до різноманітних умов.

Кожна з цих методологій має свої переваги та недоліки, і вибір конкретного підходу залежить від вимог та умов конкретної задачі.

Висновки. Інтелектуальні агенти є ключовим елементом у багатьох сучасних технологіях, що забезпечує автоматизацію складних завдань. З їх допомогою можна розв'язувати проблеми в таких сферах, як управління ресурсами, робототехніка, аналіз даних і персоналізоване обслуговування.

Математичні методи відіграють центральну роль у проектуванні, моделюванні та реалізації інтелектуальних агентів. Взаємодія між різними галузями, такими як теорія ймовірностей, оптимізація та машинне навчання, дозволяє створювати системи, здатні до адаптації, навчання та автономного функціонування. Подальший розвиток математичних основ може значно розширити можливості інтелектуальних агентів у різних сферах застосування.

Запропонований загальний алгоритм інтелектуального агента є універсальним підходом, що може бути адаптований до різних завдань. Його модульна структура забезпечує масштабованість та простоту інтеграції з іншими системами ШІ. Подальші дослідження спрямовані на вдосконалення модулів навчання та прийняття рішень, що дозволить створювати більш ефективні та автономні агенти. Майбутнє інтелектуальних агентів включає інтеграцію з новими технологіями, такими як квантові обчислення, а також розвиток етично обґрунтованих алгоритмів для забезпечення безпеки та справедливості.

Універсальна архітектура, що включає модульну структуру, динамічне навчання та адаптацію, розподілену обробку і інтеграцію різних підходів, забезпечує гнучкість, масштабованість та ефективність роботи інтелектуальних агентів у різноманітних середовищах. Додавання елементів теорії нечітких множин до математичної моделі значно покращує здатність агентів працювати з невизначеною та неточною інформацією, що підвищує точність та адаптивність прийняття рішень.

Створена математична модель дозволяє інтелектуальним агентам постійно вдосконалювати свою поведінку на основі нових нечітких даних, адаптуючись до змінних умов середовища. Модульна структура забезпечує можливість легкої заміни або оновлення окремих компонентів системи без впливу на всю архітектуру, що робить її більш гнучкою та придатною для широкого спектру застосувань.

Включення елементів нечіткості в модель дозволяє агентам ефективно працювати в умовах невизначеності, приймаючи оптимальні рішення на основі нечітких даних. Це особливо важливо для реальних застосувань, де інформація часто є неточною або неоднозначною.

Запропоновані в статті універсальна архітектура та математична модель відкривають нові можливості для розвитку інтелектуальних агентів, здатних ефективно адаптуватися до змінних умов, масштабуватися та інтегрувати різні підходи. Ці рішення сприяють підвищенню продуктивності та ефективності інтелектуальних систем, що дозволить їх широко застосовувати в різних галузях, від промисловості до обслуговування клієнтів. У кінцевому рахунку, впровадження цих рішень сприятиме створенню більш інтелектуальних, гнучких та надійних систем, здатних задовольнити потреби сучасного світу.

Напрями подальших досліджень:

1. Розробка алгоритмів для швидкої адаптації до середовищ із невідомими правилами або динамікою.
2. Моделювання систем із частковою синхронізацією між агентами.
3. Аналіз ефективності в порівнянні з існуючими рішеннями.
4. Додавання можливостей для налаштування середовища, де агент може навчатися і тестуватися.

Список використаних джерел:

1. Padgham L., Winikoff M. Developing Intelligent Agent Systems. A Practical Guide. Wiley, 2004. P. 225. DOI: 10.1002/0470861223
2. Saadi A., Maamri R., Sahnoun Z. Behavioral flexibility in Belief-Desire- Intention (BDI) architectures. *Multiagent and Grid Systems*. 2020. № 16(4). P. 343–377. DOI: <https://doi.org/10.3233/MGS-200335>
3. De Silva L., Meneguzzi F., Logan B. BDI Agent Architectures: A Survey, Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence Survey track. 2020, P. 4914–4921. DOI: <https://doi.org/10.24963/ijcai.2020/684>
4. Ekinici E., Halaç, T., Erdur C., Çetin Ö., Cakirlar I., Dikenelli O. Satisfying agent goals by executing different task semantics: HTN, OWL-S or plug one yourself. *Autonomous Agents and Multi-Agent Systems*. 2013. № 26(2). DOI: <https://doi.org/10.1007/s10458-011-9185-2>
5. Singh D., Sardina S., Padgham L., James G. Integrating Learning into a BDI Agent for Environments with Changing Dynamics. *International Joint Conference on Artificial Intelligence*. 2011. P. 2525–2530. DOI: <https://doi.org/10.5591/978-1-57735-516-8/IJCAI11-420>
6. Zhang H., i Huang S. Y., A general framework for parallel BDI agents in dynamic environments. *Web Intelligence and Agent Systems Journal*. 2008. № 6(3). P. 327–351. DOI: <https://doi.org/10.1109/IAT.2006.8>
7. Germano R., Lakhmi C. J. Intelligent Agents. Theory and Applications. Springer Berlin, Heidelberg. 2004. P. 402. DOI: <https://doi.org/10.1007/978-3-540-44401-5>
8. Khosla R., Dillon T. Engineering Intelligent Hybrid Multi-Agent Systems. Springer New York. 1997. P. 410. DOI: <https://doi.org/10.1007/978-1-4615-6223-8>
9. Bryson, J. Cross-paradigm analysis of autonomous agent architecture. *Journal of Experimental & Theoretical Artificial Intelligence*. 2000. № 12(2). P. 165–189. DOI: <https://doi.org/10.1080/095281300409829>
10. Rumbell T., Barnden J., Denham S., Wennekers T. Emotions in autonomous agents: comparative analysis of mechanisms and functions. *Auton Agent Multi-Agent Syst*. 2012. № 25. P. 1–45. DOI: <https://doi.org/10.1007/s10458-011-9166-5>

-
11. Cruz A., dos Santos A. V., Santiago R. H. N., Bedregal B. A Fuzzy Semantic for BDI Logic. *Fuzzy Information and Engineering*. 2021. № 13(2). P. 139-153. DOI: <https://doi.org/10.1080/16168658.2021.1915455>
 12. Calegari R., Ciatto G., Mascardi, V., Omicini A. Logic-based technologies for multi-agent systems: a systematic literature review. *Auton Agent Multi-Agent Syst*. 2021. № 35(1). <https://doi.org/10.1007/s10458-020-09478-3>
 13. Круглик В. С., Прокоф'єв Є. Г., Маринов, А. В. Аналіз можливостей використання інтелектуальних агентів в адаптивній системі електронного навчання. *Педагогічні науки: теорія та практика*. 2021. № 4. С. 295–302. Режим доступу: <https://doi.org/10.26661/2786-5622-2021-4-44>
 14. Лопатто І. Ю., Говорущенко Т. О., Капустян М. В. Інтелектуальний агент верифікації врахування інформації предметної галузі в процесі розроблення програмних систем. *Вісник Хмельницького національного університету*. 2022. № 1. С. 116–119. URL: <http://journals.khnu.km.ua/vestnik/?p=12131>
 15. Маринов А. В., Круглик В. С. Використання інтелектуальних програмних агентів для створення адаптивного середовища електронного навчання на базі lms moodle. *Міжнародна науково-практична конференція «Цифрова трансформація та диджитал технології для сталого розвитку всіх галузей сучасної освіти, науки і практики»*. Zbiór prac Tom 2. 2023. С. 306–308. URL: https://repo.btu.kharkov.ua/bitstream/123456789/29446/1/zbiór_prac_tom_2_26012023-306-308.pdf
 16. Noulamo T., Djimeli-Tsajio A., Kameugne R., Lienou, J. A Generic Intelligent Agent Design Approach Based on Artificial Neural Networks. *World Journal of Engineering and Technology*. 2023. № 11. С. 682–697. URL: [10.4236/wjet.2023.114046](https://doi.org/10.4236/wjet.2023.114046)
 17. Christie, S.H., Chopra, A. K., Singh, M. P. Mandrake: Multiagent Systems as a Basis for Programming Fault-Tolerant Decentralized Applications. *Autonomous Agents and Multi-Agent Systems*. 2022. № 36, A. № 16. DOI: <https://doi.org/10.1007/s10458-021-09540-8>
 18. Latham N. Types of intelligent agent, 2024. URL: <https://www.probecx.com/en-au/blog/types-of-intelligent-agent>
 19. Intelligent Agent, 2023. URL: https://www.larksuite.com/en_us/topics/ai-glossary/intelligent-agent
 20. Harjyot K. What are AI Agents: Types, Benefits, Applications, and Examples. 2024. URL: <https://www.signitysolutions.com/blog/ai-agents>
 21. Hiren Dhaduk. What is an AI Agent? Characteristics, Advantages, Challenges, Applications. 2023. URL: <https://www.simform.com/blog/ai-agent/>
 22. Крихівський М. В., Крихівська С. М. Проблеми людино-машинної взаємодії в контексті штучного інтелекту. *Збірник тез доповідей науково-практичної конференції «Інформаційні технології в освіті, техніці та промисловості»*. 2024. С. 176–177. URL: <https://stlnau.in.ua/samoosvita/item/2024/iit241010.pdf>

References:

1. Padgham, L., Winikoff, M. (2004). *Developing Intelligent Agent Systems. A Practical Guide*. Wiley, P. 225. DOI: [10.1002/0470861223](https://doi.org/10.1002/0470861223)
2. Saadi, A., Maamri, R., Sahnoun, Z. (2020). Behavioral flexibility in Belief-Desire- Intention (BDI) architectures. *Multiagent and Grid Systems*. № 16(4). P. 343–377. DOI: <https://doi.org/10.3233/MGS-200335>
3. De Silva, L. Meneguzzi, F., Logan, B. (2020). BDI Agent Architectures: A Survey, Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence Survey track. P. 4914–4921. DOI: <https://doi.org/10.24963/ijcai.2020/684>
4. Ekinici, E., Halaç, T., Erdur, C., Çetin, Ö., Cakirlar, I., Dikenelli, O. (2013). Satisfying agent goals by executing different task semantics: HTN, OWL-S or plug one yourself. *Autonomous Agents and Multi-Agent Systems*. № 26(2). DOI: <https://doi.org/10.1007/s10458-011-9185-2>
5. Singh, D., Sardina, S., Padgham, L., James, G. (2011). Integrating Learning into a BDI Agent for Environments with Changing Dynamics. *International Joint Conference on Artificial Intelligence*. P. 2525–2530. DOI: <https://doi.org/10.5591/978-1-57735-516-8/IJCAI11-420>
6. Zhang, H., Huang, S. Y. (2008). A general framework for parallel BDI agents in dynamic environments. *Web Intelligence and Agent Systems Journal*. № 6(3). P. 327–351. DOI: <https://doi.org/10.1109/IAT.2006.8>
7. Germano, R., Lakhmi, C. J. (2004). *Intelligent Agents. Theory and Applications*. Springer Berlin, Heidelberg. P. 402. DOI: <https://doi.org/10.1007/978-3-540-44401-5>
8. Khosla, R., Dillon, T. (1997). *Engineering Intelligent Hybrid Multi-Agent Systems*. Springer New York. P. 410. DOI: <https://doi.org/10.1007/978-1-4615-6223-8>
9. Bryson, J. (2000). Cross-paradigm analysis of autonomous agent architecture. *Journal of Experimental & Theoretical Artificial Intelligence*. № 12(2). P. 165–189. DOI: <https://doi.org/10.1080/095281300409829>
10. Rumbell, T., Barnden, J., Denham, S., Wennekers, T. (2012). Emotions in autonomous agents: comparative analysis of mechanisms and functions. *Auton Agent Multi-Agent Syst*. № 25. P. 1–45. DOI: <https://doi.org/10.1007/s10458-011-9166-5>
11. Cruz, A., dos Santos, A. V., Santiago, R. H. N., Bedregal, B. (2021). A Fuzzy Semantic for BDI Logic. *Fuzzy Information and Engineering*. № 13(2). P. 139–153. DOI: <https://doi.org/10.1080/16168658.2021.1915455>

-
12. Calegari, R., Ciatto, G., Mascardi, V., Omicini, A. (2021). Logic-based technologies for multi-agent systems: a systematic literature review. *Auton Agent Multi-Agent Syst.* № 35(1). <https://doi.org/10.1007/s10458-020-09478-3>
 13. Kruhlyk, V. S., Prokofiev, Ye. H., Marynov, A. V. (2021). Analiz mozhlyvostei vykorystannia intelektualnykh ahentiv v adaptivnii systemi elektronnoho navchannia [Analysis of the possibilities of using intelligent agents in an adaptive e-learning system]. *Pedahohichni nauky: teoriia ta praktyka.* № 4, S. 295–302. <https://doi.org/10.26661/2786-5622-2021-4-44>
 14. Lopatto, I. Yu., Hovorushchenko, T. O., Kapustian, M. V. (2022). Intelektualnyi ahent veryfikatsii vrakhuvannia informatsii predmetnoi haluzi v protsesi rozroblennia prohramnykh system [Intelligent agent for verification of the consideration of subject area information in the process of developing software systems]. *Visnyk Khmelnytskoho natsionalnoho universytetu*, № 1. S.116–119 Retrieved from: <http://journals.khnu.km.ua/vestnik/?p=12131>
 15. Marynov, A. V., Kruhlyk, V. S. (2023). Vykorystannia intelektualnykh prohramnykh ahentiv dlia stvorennia adaptivnoho seredovyscha elektronnoho navchannia na bazi lms moodle [Using intelligent software agents to create an adaptive e-learning environment based on lms moodle]. *Mizhnarodna naukovo-praktychna konferents «Tsyfrova transformatsiia ta dydzhytal tekhnolohii dlia staloho rozvytku vsikh haluzei suchasnoi osvity, nauky i praktyky».* Zbiór prac Tom 2. S. 306–308. Retrieved from: https://repo.btu.kharkov.ua/bitstream/123456789/29446/1/zbior_prac_tom_2_26012023-306-308.pdf
 16. Noulamo, T., Djimeli-Tsajio A., Kameugne R., Lienou, J. (2023). A Generic Intelligent Agent Design Approach Based on Artificial Neural Networks. *World Journal of Engineering and Technology.* № 11. C. 682–697. Retrieved from: 10.4236/wjet.2023.114046
 17. Christie, S.H., Chopra, A. K., Singh, M. P. (2022). Mandrake: Multiagent Systems as a Basis for Programming Fault-Tolerant Decentralized Applications. *Autonomous Agents and Multi-Agent Systems.* № 36, A. № 16. DOI: <https://doi.org/10.1007/s10458-021-09540-8>
 18. Latham N. Types of intelligent agent, 2024. Retrieved from: <https://www.probecx.com/en-au/blog/types-of-intelligent-agent>
 19. Intelligent Agent, (2023). Retrieved from: https://www.larksuite.com/en_us/topics/ai-glossary/intelligent-agent
 20. Harjyot, K. (2024). What are AI Agents: Types, Benefits, Applications, and Examples. Retrieved from: <https://www.signitysolutions.com/blog/ai-agents>
 21. Hiren Dhaduk. (2023). What is an AI Agent? Characteristics, Advantages, Challenges, Applications. Retrieved from: <https://www.simform.com/blog/ai-agent/>
 22. Krykhivskyi, M. V., Krykhivska, S. M. (2024). Problemy liudyno-mashynnoi vzaiemodii v konteksti shuchnoho intelektu [Problems of human-machine interaction in the context of artificial intelligence]. *Zbirnyk tez dopovidei naukovo-praktychnoi konferentsii «Informatsiini tekhnolohii v osviti, tekhnitsi ta promyslovosti».* S. 176–177. Retrieved from: <https://stlnau.in.ua/samoosvita/item/2024/iit241010.pdf>

Дата надходження статті: 24.10.2025

Дата прийняття статті: 17.11.2025

Опубліковано: 30.12.2025