

Міністерство освіти і науки України
Національний університет «Києво-Могилянська академія»
Факультет інформатики
Кафедра математики

Магістерська робота

освітній ступінь – магістр

На тему: **«МОДЕЛЮВАННЯ ЕКОНОМІЧНИХ ПРОЦЕСІВ ЗА ДОПОМОГОЮ
КЕРОВАНИХ МАРКОВСЬКИХ ПОЛІВ»**

Виконав: студент 2-го року навчання,
Спеціальності

113 Прикладна математика

Магдич Назар Юрійович

Керівник Чорней Р. К.

кандидат фізико-математичних наук,
доцент

Рецензент _____
(прізвище та ініціали)

Магістерська робота захищена

З оцінкою _____

Секретар ЕК _____

« ____ » _____ 20 ____ р.

Київ-2025

Національний університет "Києво-Могилянська академія"

Факультет інформатики

Кафедра математики

Освітній ступінь магістра

Спеціальність 113 Прикладна математика

Освітньо-наукова програма "Прикладна математика"

ЗАТВЕРДЖУЮ

Завідувач кафедри _____

_____ " ____ " _____ 20 ____ року

ЗАВДАННЯ

ДЛЯ МАГІСТЕРСЬКОЇ РОБОТИ СТУДЕНТУ

Магдичу Назару Юрійовичу

1. Тема роботи: "Моделювання економічних процесів за допомогою керованих марковських полів"

керівник роботи: Чорней Руслан Костянтинович, кандидат фіз.-мат. наук, доцент
затверджені наказом вищого навчального закладу

від " ____ " _____ 20 ____ року № _____

2. Строк подання студентом роботи:

3. План роботи:

1. Вступ
2. Теоретична частина
3. Моделювання економічного процесу
4. Приклад розв'язання моделі
5. Висновки
6. Список літератури

Графік підготовки магістерської роботи до захисту

№ з/п	ПЕРЕЛІК РОБІТ	Термін виконання	Дата ознайомлення наукового керівника	Підпис наукового керівника	Примітки
1.	Вибір теми, затвердження її на засіданні кафедри та закріплення наукового керівника Узгодження календарного графіка підготовки магістерської роботи. Ознайомлення студента з критеріями оцінювання магістерської роботи	жовтень	10.10.2025		
2.	Вивчення джерел, літератури, періодичних видань, наукових публікацій, збір та узагальнення фактів, даних	жовтень-листопад	15.11.2025		
3.	Складання плану магістерської роботи та узгодження із науковим керівником	листопад	22.11.2025		
4.	Постановка експерименту, аналіз отриманих результатів наукового дослідження	листопад-березень	19.03.2025		
5.	Проміжний контроль виконання роботи	лютий	28.02.2025		
6.	Написання кваліфікаційної роботи в цілому, ознайомлення з її першим ва-ріантом наукового керівника.	січень-березень	29.03.2025		
7.	Повне завершення написання кваліфікаційної роботи, оформлення її згідно з вимогами й подання на відгук науковому керівнику	квітень-травень	07.05.2025		
8.	Доповідь на XIII Всеукраїнській науковій конференції молодих математиків	травень	09.05.2025		
9.	Подання магістерської роботи для перевірки письмових робіт студентів На-УКМА на відповідність вимогам академічної доброчесності.	початок червня	06.06.2025		
10.	Подання на зовнішню рецензію	початок червня	06.06.2025		
11.	Підготовка до захисту магістерської роботи на засіданні кафедри: написання доповіді та виготовлення ілюстративного матеріалу	до 6 червня	06.06.2025		
12.	Підготовка до захисту магістерської роботи на засіданні кафедри: написання доповіді та виготовлення ілюстративного матеріалу	до 6 червня	06.06.2025		
13.	Подання магістерської роботи на кафедру з усіма супроводжуючими документами.	до 6 червня	06.06.2025		
14.	Публічний захист перед екзаменаційною комісією	Згідно з розкладом роботи ЕК	13.06.2025		

Графік узгоджено "___" _____ 20__ р.

Науковий керівник _____ (ПІБ)

Виконавець кваліфікаційної роботи _____ (ПІБ)

ЗМІСТ

АНОТАЦІЯ.....	3
ВСТУП.....	4
1. ОСНОВНІ ТЕОРЕТИЧНІ ВІДОМОСТІ.....	6
2. МОДЕЛЮВАННЯ ЕКОНОМІЧНОГО ПРОЦЕСУ.....	12
2.1 Етапи побудови моделі за допомогою керованих марковських полів	12
2.2 Моделі з компактними множинами станів та дій.	14
2.3 Процес пошуку оптимальних стратегій.....	16
3. РОЗВ'ЯЗАННЯ ЕКОНОМІЧНОЇ ЗАДАЧІ НА СКІНЧЕНОМУ	
ГОРИЗОНТІ	22
ВИСНОВКИ	27
СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ	28
ДАДАТОК А	29
ДАДАТОК В	30

АНОТАЦІЯ

У роботі розглядається задача пошуку оптимальних стратегій керування марковськими полями на скінченному часовому горизонті. Актуальність дослідження зумовлена потребою у моделюванні складних динамічних систем з урахуванням випадкових впливів, зокрема в економіці. Наведено теоретичні відомості про керовані марковські поля з локальною взаємодією. Запропоновано алгоритм пошуку оптимальних стратегій на скінченному горизонті. Результати можуть бути використані в задачах економічного прогнозування та управління ресурсами.

Ключові слова: випадкові поля, оптимальне керування, математична модель, локальні стратегії, стохастичний процес, економічні процеси, марковська властивість.

ABSTRACT

The problem of finding optimal control strategies for Markov fields on a finite time horizon is considered in this paper. The relevance of the study is conditioned by the need of modeling complex dynamic systems influenced with random effects, including economic processes. Theoretical information about controlled Markov fields with local interaction is presented. An algorithm for finding optimal strategies on a finite horizon is proposed. The results can be used in problems of economic forecasting and resource management.

Keywords: random fields, optimal control, mathematical model, local strategies, stochastic process, economic processes, Markov property.

ВСТУП

Актуальність теми: Керовані поля Маркова є інструментом моделювання систем, що змінюються в часі та просторі. Вони дозволяють врахувати вплив випадкових процесів на систему та керування, що здійснюється агентами. Спираючись на такі моделі можливо прогнозувати та оптимізовувати процеси управління та прийняття рішень в умовах невизначеності та ризику. Подібні моделі можуть бути корисними для розв'язання задач в області управління ресурсами, економічного прогнозування, логістики. Вивчення керованих марковських полів та розробка нових методів їх застосування є актуальними через зростаючу потребу в інструментах аналізу даних. В роботі [1] сформульовано та доведено теорему про існування та єдиність оптимальної стратегії для керованого марковського процесу з локальною взаємодією, в роботі [2] описано алгоритм, який можна застосувати для розв'язання такої моделі. Основним недоліком алгоритму є необхідність визначення ядра переходу, що є трудомістким процесом. Проте, для багатьох економічних процесів немає необхідності оптимізовувати стратегію на нескінченному горизонті, але є потреба віднайти оптимальну стратегію керування на скінченному проміжку часу.

Об'єкт дослідження: Керовані марковські поля.

Предмет дослідження: Оптимальні стратегії керування марковськими полями.

Мета дослідження: Розробити алгоритм пошуку оптимальних марковських стратегій на скінченному горизонті для вирішення прикладних задач з економіки.

Завдання:

- Дослідити теоретичні відомості про керовані марковські поля та їх властивості.
- Визначити області застосування керованих марковських полів в економіці та проаналізувати можливість їх застосування для вирішення практичних задач.

- Описати процес побудови моделі та пошуку оптимальних стратегій керування на скінченному горизонті.
- Провести експериментальні дослідження для підтвердження дієвості розроблених методів.

Методи дослідження: У роботі використовуються методи математичного моделювання та чисельні методи. Алгоритми реалізовано за допомогою мови програмування python.

Наукова новизна: Розробка алгоритму пошуку оптимальних стратегій керування марковськими полями на скінченному горизонті.

Практичне значення: Можливість застосування розробленого алгоритму для розв'язання прикладних задач з економіки.

Результати роботи були подані в доповіді на XIII Всеукраїнській науковій конференції молодих математиків. (Додаток А) та опубліковані в збірці тез конференції [3].

Структура роботи: Робота складається з вступу, теоретичної частини, опису процесу побудови моделі, розв'язання економічної задачі, висновків та списку літератури. У вступі розкрито актуальність теми, об'єкт, предмет, мету, завдання, методи дослідження, наукову новизну та практичне значення роботи. В теоретичній частині описані основні властивості та теоретичні основи керованих марковських полів з локальною взаємодією, описана задача з пошуку оптимальної стратегії керування. Розробка моделей містить опис процесу побудови моделей економічних процесів на основі керованих марковських полів, їх переваги та обмеження. подано означення оптимальної стратегії на скінченному горизонті, запропоновано та обґрунтовано алгоритм пошуку такої стратегії. Експериментальна частина присвячена розв'язанню задачі за допомогою програмних засобів. У висновках підведені підсумки роботи та окреслені перспективи подальших досліджень.

1. ОСНОВНІ ТЕОРЕТИЧНІ ВІДОМОСТІ

У цьому розділі наведені визначення, необхідні для формалізації та побудови моделей просторових систем з локальним керуванням та дискретним часом. Зокрема, випадкові марковські поля з локальною взаємодією компонентів, де кожен елемент системи може обмінюватися інформацією або впливати лише на сусідні елементи.

Керування представлене у вигляді локальних детермінованих стратегій, які є функціями від поточного стану повного околу агента. Це означає, що рішення приймається на основі інформації, доступної лише в межах повного околу, без глобального узгодження або централізованого планування.

Подано умовні позначення, які використовуються протягом усієї роботи. Вміст розділу спирається на теоретичні відомості з роботи [1].

Розглянемо скінченний, ненаправлений $G = (V, B)$ граф, у якому

V – множина вершин.

$B = \{(k, j)\}$ – множина ребер, що з'єднують вершини k та j .

Для кожної вершини $k \in V$ визначено

$N(k) = \{j: \{(k, j)\} \in B\}$ – окіл вершини k .

$\tilde{N}(k) = N(k) \cup \{k\}$ – повний окіл вершини k .

$X_k = \{x_k\}$ – множина можливих станів.

$X = \times_{k \in V} X_k$ – глобальна множина станів системи.

$X_k = \times_{j \in \tilde{N}(k)} X_j$ – множина станів повного околу вершини k .

$(\Omega, \mathcal{F}, P)$ – ймовірнісний простір.

Означення 1.1 Дискретне марковське випадкове поле - випадкова величина ξ , що визначена на просторі $(\Omega, \mathcal{F}, P)$, приймає значення з X , та задовольняє умову:

$$P(\xi_k = x_k | \xi_{V-\{k\}} = x_{V-\{k\}}) = P(\xi_k = x_k | \xi_{N(k)} = x_{N(k)}).$$

Для моделювання економічних процесів доцільно вважати, що значення ξ_k залежить не лише стану околу $N(k)$ вершини k , але від повного її околу $\tilde{N}(k)$. Також, варто вказати залежність стану системи від дискретного часу t , $t \in \mathbb{N}$. Таким чином визначимо марковський процес з локальною, синхронною взаємодією компонентів на графі (V, B) .

Означення 1.2 Марковський процес з локальною, синхронною взаємодією – випадкова величина ξ_k^t , що приймає значеннями з X_k визначена на $(\Omega, \mathcal{F}, P)$ просторі, що залежить від часу t , має локальні ймовірності переходу:

$$P(\xi_k^{t+1} = y_k | \xi_{V-\{k\}}^t = x_{V-\{k\}}^t, \dots, \xi_{V-\{k\}}^0 = x_{V-\{k\}}^0) = P(\xi_k^{t+1} = y_k | \xi_{\tilde{N}(k)}^t = x_{\tilde{N}(k)}^t).$$

Та має синхронні ймовірності переходу:

$$P(\xi_K^{t+1} = x_K^{t+1} | \xi^t = x^t) = \prod_{k \in K} P(\xi_k^{t+1} = x_k^{t+1} | \xi^t = x^t), \quad K \subset V, \quad x^t, x^{t+1} \in X.$$

В подальшому визначений таким чином процес називатиметься випадковим марковським полем.

Введемо керування для випадкового марковського поля: Нехай у кожній вершині $k \in V$, в кожний момент часу $t \in \mathbb{N}$ незалежний агент приймає рішення з обмеженої, скінченної множини A_k .

A_k – множина можливих дій в вершині $k \in V$.

$A = \times_{k \in V} A_k$ – глобальний простір можливих дій в системі.

$\Delta_k^t(x^0, \dots, x^t)$ – рішення прийняте в момент часу t у вершині k , що залежить від всієї історії системи.

$\Delta^t(x^0, \dots, x^t) = \{\Delta_k^t(x^0, \dots, x^t)\}$ – рішення прийняте в момент часу t .

Визначимо локальну стратегію керування:

Означення 1.3 Локальна стратегія δ - множина $\{\delta_k, k \in V\}$, у якій δ_k - послідовність $\{\Delta_k^t, k \in V, t \in \mathbb{N}\}$ таких, що $\Delta_k^t(x_{\tilde{N}(k)}^0, \dots, x_{\tilde{N}(k)}^t)$ – відображення $X_{\tilde{N}(k)}^0 \times \dots \times X_{\tilde{N}(k)}^t \rightarrow A_k$, тобто рішення, що приймається на основі всієї історії станів повного околу вершини k до моменту часу t .

У цій роботі розглядається специфічний клас стратегій керування, а саме марковські локальні стаціонарні детерміновані стратегії, які застосовуються до просторових систем із дискретним часом. Марковська властивість для стратегій передбачає, що вибір дії не залежить від попередньої історії, а лише від поточного стану елементів повного околу на момент часу t . Стаціонарність передбачає незмінні в часі $\Delta_k: X_{\tilde{N}(k)} \rightarrow A_k$, що формують стратегію δ_k . Таким чином, локальна стратегія керування є функцією від повного околу вершини.

Означення 1.4 Марковська локальна стратегія $\delta = \{\delta_k, k \in V\}$ – локальна стратегія, для якої виконується умова Маркова: $\Delta_k^t(x_{\tilde{N}(k)}^0, \dots, x_{\tilde{N}(k)}^t) = \Delta_k^t(x_{\tilde{N}(k)}^t)$, Тобто рішення залежить лише від стану повного околу в поточний момент часу t .

Означення 1.5 Марковська локальна стаціонарна стратегія $\delta = \{\delta_k, k \in V\}$ – локальна марковська стратегія, для якої виконується умова: $\Delta_k^{t'}(x_{\tilde{N}(k)}) = \Delta_k^{t''}(x_{\tilde{N}(k)})$, Тоді така стратегія визначається функціями $\Delta_k: X_{\tilde{N}(k)} \rightarrow A_k, k \in V$.

Означення 1.6 Локальна допустима стратегія δ – локальна стратегія, для якої

1) $A_k^t(x_{\tilde{N}(k)})$ - множина допустимих дій при $\xi_{\tilde{N}(k)}^t = x_{\tilde{N}(k)}, \forall k \in V, \forall x_{\tilde{N}(k)} \in X_k$.

$$2) \Delta^t(x^0, \dots, x^t) = \{\Delta_k^t(x_{\tilde{N}(k)}^0, \dots, x_{\tilde{N}(k)}^t)\} \in A^t(x^t).$$

Множину всіх допустимих детермінованих локальних стратегій позначимо, як LD .

Множину всіх допустимих детермінованих локальних марковських стратегій позначимо, як LD_S .

Означення 1.7 Керований процес з локальною взаємодією на графі $G = (V, B)$ – пара (ξ, δ) , де $\xi = (\xi^t, t \in \mathbb{N})$ – стохастичний процес, визначений в 1.2 з множиною станів $X = \times_{k \in V} X_k$. $\delta = \{\delta_k, k \in V\}$ - локальна стратегія. Визначені ймовірності переходу ξ визначені наступним чином

$$\begin{aligned} P\{\xi_S^{t+1} = x_S | \xi^0 = x^0, \Delta^0(\xi^0) = a^0, \dots, \xi^t = x^t, \Delta^t(\xi^t) = a^t\} &= \\ &= P\{\xi_S^{t+1} = x_S | \xi^0 = x^0, \dots, \xi^t = y, \Delta^t(\xi^0, \dots, \xi^t) = a\} = \\ &= P\{\xi_S^{t+1} = x_S | \xi^t = y, \Delta^t(\xi^0, \dots, \xi^t) = a\} = \\ &= \prod_{j \in S} P\{\xi_j^{t+1} = x_j | \xi^t = y, \Delta^t(\xi^0, \dots, \xi^t) = a\} = \\ &= \prod_{j \in S} P\{\xi_j^{t+1} = x_j | \xi_{\tilde{N}(j)}^t = y_{\tilde{N}(j)}, \Delta_j^t(\xi_{\tilde{N}(j)}^0, \dots, \xi_{\tilde{N}(j)}^t) = a_j\} = \\ &= \prod_{j \in S} Q_j\{x_j, y_{\tilde{N}(j)}, a_j\} \cdot 1_{\{a\}}(\Delta^t(\xi^0, \dots, \xi^t)) = \\ &= Q_S(x_S, y, a) \cdot 1_{\{a\}}(\Delta^t(\xi^0, \dots, \xi^t)), \quad S \subseteq V, \quad y \in X, \quad a \in A(y). \end{aligned}$$

для кожної вершини графа k визначено локальні ядра переходу:

$Q_k = (x_k | y_{\tilde{N}(k)}, a_k)$, що визначають незалежний від часу закон переходу всієї системи: $Q = (x | y, a) = P(\xi^{t+1} = x_k^{t+1} | \xi^t = y, \Delta^t = a)$.

Задача пошуку оптимальної стратегії вимагає визначення цільової функції. В залежності від сфери застосування моделі це може бути функція втрат чи винагород, що включатимуть в себе втрати, або винагороди отримані в кожен момент часу t .

Означення 1.8 Однокрокова втрата(винагорода) системи в момент часу $t \in \mathbb{N}$ - $r(x^t, a^t) \geq 0$, що залежить від стану x^t та дії a^t .

Означення 1.9 Очікувані середні втрати(винагорода) $Q_T^\delta(y)$ при використанні стратегії δ , та початковому стані системи $\xi^0 = y$ визначається так

$$Q_T^\delta(y) = E_y^\delta \frac{1}{1+T} \sum_{t=0}^T r(\xi^t, \Delta^t(\xi^0, \dots, \xi^t)).$$

Математичне сподівання E_y^δ відповідає керованому процесу (ξ, δ) . Ціль полягає в знаходженні такої стратегії δ , що максимізує очікувану винагороду на нескінченному горизонті $T \rightarrow \infty$.

Означення 1.10 Цільова функція за умови прийняття стратегії δ та початкового стану $y \in X - R_y^\delta$ така, що

$$R_y^\delta = \limsup_{T \rightarrow \infty} (Q_T^\delta(y)), \quad Q_T^\delta - \text{очікувані середні втрата,}$$

або

$$R_y^\delta = \limsup_{T \rightarrow \infty} (Q_T^\delta(y)), \quad Q_T^\delta - \text{очікувані середня винагорода.}$$

Означення 1.11 Оптимальна стратегія керування δ - елемент множини допустимих марковських стаціонарних локальних стратегій, що задовольняє умови:

$$p(y) = \inf_{\delta \in LD} (\limsup_{T \rightarrow \infty} (Q_T^\delta(y))), \quad \forall y \in X,$$

$$Q_T^\delta - \text{очікувані середні втрати,}$$

або:

$$p(y) = \sup_{\delta \in LD} (\liminf_{T \rightarrow \infty} (Q_T^\delta(y))), \quad \forall y \in X,$$

Q_T^δ — очікувана середня винагорода,

де $p(y)$ — значення, яке набуває оптимальна стратегія.

Ключовою для цієї роботи є теорема про існування та єдиність оптимальної стратегії з множини допустимих детермінованих стаціонарних локальних марковських стратегій, доведення наведено в роботі [1].

Теорема 1.1 Для керованого процесу (ξ, δ) на графі $G = (V, B)$, зі скінченною множиною станів X та скінченною, незалежною від часу t множиною дій A в множині допустимих детермінованих локальних стратегій LD існує єдина оптимальна стратегія, що належить множині LD_S .

2. МОДЕЛЮВАННЯ ЕКОНОМІЧНОГО ПРОЦЕСУ

2.1 Етапи побудови моделі за допомогою керованих марковських полів.

Для побудови моделі економічного процесу за допомогою керованих марковських полів необхідно виконати наступні кроки.

- Формалізувати структуру системи у вигляді неорієнтованого графа $G = (V, B)$. Множина вершин V відображає набір економічних агентів, (підприємства, споживачі, регіони, фірми). Кожен агент здатен приймати рішення, що впливають на власний стан на стани пов'язаних ребром вершин. Множина ребер B відображає взаємозв'язок між агентами, що може включати територіальне сусідство, інфраструктурний зв'язок, конкуренцію, кооперацію або іншу взаємодію, яку необхідно передбачити в моделі. Важливо дотримуватись однорідності зв'язків, відображених ребрами в рамках однієї моделі, тобто інтерпретація взаємодії між будь-якими двома агентами має бути єдиною. Процес побудови графа G вимагає ретельного дослідження економічного процесу для визначення ключових вершин та ребер.

- Для кожної з вершин-агентів визначити множини станів X_k , та дій A_k , $k \in V$. Стан вершини є формалізованим відображенням певної характеристики або сукупності характеристик що залежать від стану повного околу вершини та дії, яку приймає агент в момент часу t . Дія агента у вершині k – елемент A_k , що підпорядковується стратегії δ , через який агент може впливати на стан повного околу шляхом зміни ймовірності переходу між станами.

- Визначити цільову функцію.

Залежно від характеру моделі та економічної мети, яку вона відображає, необхідно віднайти стратегію, що мінімізує очікувані середні витрати, або максимізує очікувану середню винагороду при $T \rightarrow \infty$. Цільова функція задає оптимізаційну задачу, яка вирішується або для окремого агента, або для системи в цілому, і визначає

оптимальну локальну стратегію керування. Економічним процесам притаманні невизначеність та ризики, що нарастають зі збільшенням горизонту планування. Ці фактори варто враховувати в моделях. Для цього до функції винагород вводять дисконт-коефіцієнт $\theta \in (0,1]$, що зменшуватиме цінність віддалених в часі винагород, та надаватиме перевагу наближеним винагородам. Тоді цільова функція має вигляд:

$$R_y^\delta = \limsup_{T \rightarrow \infty} \left(E_y^\delta \frac{1}{1 + T} \sum_{t=0}^T \theta \cdot r(\xi^t, \Delta^t) \right).$$

Для застосування методів обчислення оптимальної стратегії для моделі керованого марковського процесу (ξ, δ) , запропонованих в [2] необхідно переконатись, що модель економічного процесу задовольняє ряд умов:

- Множина станів X_k та дій A_k скінченні та обмежені $\forall k \in V$.

Обмеженість множин є природньою для більшості економічних процесів, оскільки можливі інтерпретації стану можуть включати запас продукції, рівень доходу, зайнятість, борг або виробничу потужність – усі ці характеристики не можуть набувати нескінченних значень. Також, множина дій обмежена спроможностями агента. Скінченність множин зумовлена необхідністю обрахування ймовірностей переходу та пошуку оптимальної стратегії. Деякі міркування щодо побудови та розв'язання моделей, де стани та дії агентів є проміжками на числовій прямій наведено в підрозділі 2.2 цього розділу.

- ξ – Марковський процес з локальною, синхронною взаємодією.

Перехід між станами для всіх агентів відбувається синхронно в дискретному часі t . Це дозволяє моделювати процеси для яких притаманна сезонність, або стрибкоподібні зміни. Більш того, в цьому випадку, деякі елементи моделі можливо представити у вигляді стохастичних процесів, наприклад, при моделюванні поведінки агентів на фінансовому ринку можливо представити коливання ціни активів у вигляді Броунівського руху. Марковська властивість виконується для

економічних процесів, у яких інформація про поточний стан є достатньою для прийняття рішення, в іншому випадку, модель не може бути застосована, і необхідно шукати інший інструмент, наприклад, моделі з пам'яттю.

- δ – марковська локальна стаціонарна детермінована стратегія.

Це обмеження обумовлено результатами роботи [1], де доведено існування і єдиність оптимальної стратегії в множині допустимих локальних стаціонарних, детермінованих стратегій.

2.2 Моделі з компактними множинами станів та дій.

Деякі характеристики економічних процесів, які відповідають стану системи можуть бути представлені у вигляді проміжків на числовій прямій. Простим підходом до побудови моделей таких процесів може бути наближення проміжків дискретними множинами шляхом поділу проміжків на однакові відрізки. розглянемо приклад задачі з визначення оптимальної цінової політики:

Нехай декілька фірм конкурують на ринку та змушені підлаштовувати свою цінову політику, щоб максимізувати прибуток. При встановленні ціни, фірма має враховувати попит, що залежить від цін конкурентів та випадкових коливань. Задача кожного з агентів – визначити оптимальну цінову політику.

Кожна з фірм представлена, як вершина k з множини V , Ребро (k, j) між вершинами k та j відображає конкуренцію. Для кожної з вершин визначені наступні множини:

Множина станів: $X_k = d_k = [0, d_{max}]$ – попит на товар,

Множина можливих дій фірми: $A_k = p_k = [p_{min}, p_{max}]$

p_{min} – собівартість товару.

Дискретні відповідники множин надалі позначатимемо символом *. Поділимо проміжки X_k, A_k на однакові відрізки довжини $\Delta d, \Delta p$ відповідно. Елементами дискретних множин X_k^* та A_k^* стануть крайні точки відрізків. Таким чином, отримали нові множини.

Множина станів: $X_k^* = d_k^* = \{0, 0 + \Delta d, \dots, d_{max} - \Delta d, d_{max}\}$.

Множина можливих дій фірми: $A_k^* = \{p_{min}, p_{min} + \Delta p, \dots, p_{max} - \Delta p, p_{max}\}$.

Отже, можливо побудувати модель за допомогою керованих полів маркова і віднайти для неї оптимальну стратегію δ^* . Тим не менш, цей підхід має суттєвий недолік. В загальному випадку стратегія δ^* не обов'язково є наближенням шуканої оптимальної стратегії для економічного процесу з множинами-проміжками. Тобто, такий підхід не призводить до бажаного результату.

Нехай, функція втрат системи $r(x, a)$ є відокремлюваною [1] і задається наступним чином.

$$r(x, a) = \sum_{k \in V} r_k(x_{\tilde{N}(k)}, a_k).$$

Розглянемо множину дій $A_k = [a_0, a_{max}]$ деякої вершини k . В результаті поділу її на 10 рівних проміжків, отримано нову дискретну множину дій, що складається з крайніх точок проміжків: $A_k^* = \{a_0, a_1, \dots, a_{10}\}$. Нехай однокрокова втрата для вершини k : $r_k(x_{\tilde{N}(k)}, a_k)$, визначена на проміжку $[a_0, a_{max}]$ та однокрокова втрата $r_k^*(x_{\tilde{N}(k)}, a_k)$, визначена на дискретній множині A_k^* , зображені на рисунку 2.1.

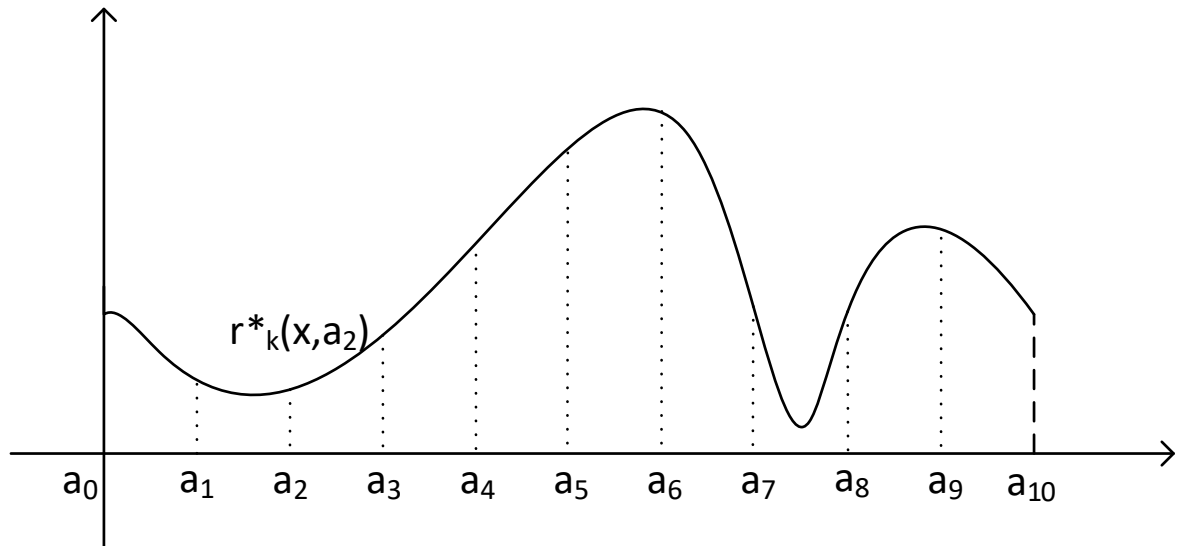


Рисунок 2.1 функції однокрокових втрат.

Як показано на рисунку 1, мінімальна втрата на неперервному проміжку досягається між станами x_7, x_8 . При цьому, функція $r_k^*(x_{\tilde{N}(k)}, a_k)$, визначена на множині A_k^* набуває свого найменшого значення в точці a_2 .

В цьому випадку вирішенням проблеми було б зменшення проміжку Δa , при цьому необхідно провести додатковий аналіз функції $r(x, a)$ для визначення достатньо малих Δa та Δx . Такий підхід є трудомістким, оскільки $r(x, a)$ – функція, що має $2 \cdot |V|$ аргументів у випадку повного графа G . Також, при зменшенні Δa та Δx , а значить – збільшенні потужностей множин X_k^* та A_k^* зростає кількість обчислень, що ускладнює пошук оптимальної стратегії δ^* . Більш детально про це в підрозділі 2.3 цього розділу.

2.3 Процес пошуку оптимальних стратегій.

Методи знаходження оптимальних стратегій керування описані в роботі [2], є складними в реалізації при великих $|X|, |A|$, оскільки вимагають визначення ядра переходу Q , що в свою чергу потребує визначення значної кількості ймовірностей

переходів, в загальному випадку $|X \times X \times A|$. Одним із способів вирішення цієї проблеми є зменшення розмірностей множин станів та дій шляхом їх агрегації, тобто об'єднання близьких елементів множин. А також, видалення зв'язків, які є незначущими для системи. Приклад застосування цих методик є в роботі [4], та узагальнено в [5]. Також метод можливо застосовувати у випадку розрідженої матриці Q , коли більшість ймовірностей переходів дорівнюють нулю або близькі до нуля. Тоді матрицю можливо зберігати у розрідженому форматі, та використовувати лише ненульові елементи, що значно зменшує кількість обчислень.

Варто розглянути моделювання економічних процесів, що обмежені в часі та мають скінченний часовий горизонт T_0 , або ж для яких змінна t відображає значні проміжки часу (рік, місяць), оскільки для багатьох процесів в економіці неможливо передбачити, чи зберігатиметься адекватність моделі через деякий значний проміжок часу T_0 . В такому разі, задача зводиться до пошуку стратегії, що оптимізує $Q_T^\delta(y)$ при $T = T_0$. Процес пошуку стратегії δ полягає в обчисленні очікуваних витрат $Q_{T_0}^{\delta'}(y)$ та пошуку мінімального значення. Значення очікуваних витрат на скінченному горизонті T_0 , за умови прийняття стратегії оптимальної δ позначатиметься: $p_{T_0}(y)$.

Означення 2.1 Стратегія δ є оптимальною на скінченному горизонті T_0 в цілому, якщо:

$$p_{T_0}(y) = \inf_{\delta \in LD_S} (Q_{T_0}^\delta(y)), \quad \forall y \in X,$$

$Q_{T_0}^\delta(y)$ – очікувані середні витрати.

Або:

$$p_{T_0}(y) = \sup_{\delta \in LD_S} (Q_{T_0}^\delta(y)), \quad \forall y \in X,$$

$Q_{T_0}^\delta(y)$ – очікувана середня винагорода.

В загальному випадку, оптимальна на скінченному горизонті стратегія δ не обов'язково існує, оскільки на скінченному горизонті очікувані середні витрати

залежать від початкового стану, тобто При різних початкових станах $y \in X$, різні стратегії δ можуть призводити до найменших втрат.

Означення 2.2 Стратегія δ_y є оптимальною на скінченному горизонті T_0 за заданого початкового стану $y \in X$ якщо:

$$p_{T_0}(y) = \inf_{\delta \in LD_S} (Q_{T_0}^\delta(y)), \quad Q_{T_0}^\delta(y) - \text{очікувані середні витрати.}$$

Або:

$$p_{T_0}(y) = \sup_{\delta \in LD_S} (Q_{T_0}^\delta(y)), \quad Q_{T_0}^\delta(y) - \text{очікувана середня винагорода.}$$

Пошук оптимальної стратегії на скінченному горизонті можна виконати алгоритмічно за визначеної структури графа, множин X_k та A_k , та однокрокових втрат $r(x, a)$.

Алгоритм 2.1 Пошук оптимальних локальних стратегій δ на скінченному горизонті T за допомогою методу оберненої математичної індукції.

- 1) Обрахувати однокрокову винагороду на останньому кроці $T: r^T(x^T, a^T)$ для усіх можливих станів системи $x^T \in X$.
- 2) Обрахувати суму однокрокової винагорода на попередньому кроці $T - 1$ та очікуваної винагорода на кроці T :

$$r^{T-1}(x^{T-1}, a^{T-1}) + E_{x^{T-1}} r^T(x^T, a^T).$$

Математичне очікування E відповідає випадковій величині ξ^t на останньому кроці $t = T$. Визначити дії $\{a_k^{t-1}, k \in V\}$, що призводять до найбільшої очікуваної винагорода за два останні кроки.

$$\sup_{a^{T-1} \in A} (r^{T-1}(x^{T-1}, a^{T-1}) + E_{x^{T-1}} r^T(x^T, a^T)).$$

- 3) В зворотньому порядку обраховувати очікувану винагороду, від кожного попереднього кроку $t = T - i$, $i = 1, 2, \dots, T$, до останнього кроку T . На кожній ітерації визначати дії $\{a_k^j, k \in V\}$ такі, що призводять до найбільшої очікуваної винагорода.

$$\sup_{\{a^{T-i}, \dots, a^{T_0}\} \in A^{T-i} \times \dots \times A^T} (r^{T-i}(x^{T-i}, a^{T-i}) + \sum_{j=T-i+1}^T E_{x^{T-i}} r^j(x^j, a^j)).$$

4) Пункт 3) алгоритму повторюється допоки $t \geq 0$. Таким чином, коли алгоритм досягає кроку $t = 0$, для кожного початкового стану $y \in X$ отримуємо послідовність дій в системі $\{a^t, t = 0, 1, \dots, T\}$ що призводять до найбільшого очікуваного виграшу при початковому стані y .

$$\sup_{a^t \in A^0 \times \dots \times A^T} \left(r^0(y, a^0) + \sum_{t=1}^T E_y r^t(x^t, a^t) \right). \quad (2.1)$$

Оскільки за означенням 1.7, керованого марковського процесу з локальною взаємодією, ймовірності переходу системи є незалежними від часової змінної t , то послідовності дій $\{a^t\}$ можливо представити у вигляді стратегій δ_y оптимальних за означенням 2.2.

Дієвість алгоритму спирається на означення 1.9 очікуваних середніх винагород. Визначимо очікувані середні витрати на скінченному горизонті T для стратегії δ_y , знайденої за допомогою алгоритму 2.1. З виразу (1) маємо:

$$\begin{aligned} \frac{\sup_{\delta_y \in LD_S} (r^0(y, a^0) + \sum_{t=1}^T E_y r^t(x^t, a^t))}{T + 1} &= \sup_{\delta_y \in LD_S} \left(\frac{1}{T + 1} \sum_{t=0}^T E_y^\delta r^t(x^t, a^t) \right) = \\ &= \sup_{\delta_y \in LD_S} \left(\frac{1}{T + 1} E_y^\delta \sum_{t=0}^T r^t(x^t, a^t) \right) = \sup_{\delta_y \in LD_S} (Q_T^\delta(y)). \end{aligned}$$

Тому знайдена за допомогою алгоритму 2.1 стратегія δ_y є оптимальною на скінченному горизонті T за початкового стану y .

Головним недоліком такого підходу є складність обчислень, що залежить від кількості вершин, структури графа та потужностей множин станів та дій. Для оцінки складності обчислень в найгіршому випадку граф $G = (V, B)$ вважаємо повним, тоді

$X_{\tilde{N}(k)} = X, \forall k \in V$; Потужності множин дій $|A_k|$ та станів $|X_k|$ у кожній з вершин вважатимемо рівними $\max_{k \in V} |A_k|, \max_{k \in V} |X_k|$ відповідно.

Проведено оцінку складності алгоритму 2.1 на повному графі G .

На першому кроці алгоритму необхідно обчислити однокрокову винагороду $r(x, a)$ для кожного можливого стану системи, тобто $|X|$ разів. А також – ймовірності переходу між станами системи, яких в найгіршому випадку може бути $|X|^2 |A|$.

На подальших кроках необхідно обчислити очікувану винагороду $E_y r(x, a)$: для кожного можливого стану y та дії a в системі. Кількість пар $|(x, a)| = |X| |A| = (\prod_{k \in V} |X_k| |A_k|)$. В найгіршому випадку кількість необхідних обчислень функції $r(x, a)$.

$$|X| \cdot |A|.$$

Однокрокова винагорода обчислюється в зворотньому порядку для кожного моменту часу $t = T - 1, \dots, 0$, тому загальна кількість обчислень:

$$O(|X| + |X|^2 |A| + T(|X| |A|)).$$

Якщо в результаті застосування алгоритму 2.1 отримали множину стратегій $\{\delta_y, y \in X\}$, таких, що $\delta_y = \delta_{y'}, \forall y, y' \in X$, то стратегія δ_y є оптимальною на скінченному горизонті за означенням 2.1, тобто є оптимальною за довільного початкового стану.

В загальному випадку, знайдена таким чином стратегія δ не обов'язково є оптимальною за означенням 1.11, але на практиці пошук такої стратегії дозволяє оптимізувати процес прийняття рішень агентів на скінченному проміжку часу.

Твердження 2.1 Оптимальна на скінченному горизонті T_0 стратегія δ . є оптимальною за означенням 1.11, якщо.

$$r(\xi^t, \Delta^t(\xi^t)) \leq r(\xi^t, \Delta'^t(\xi^t)), \quad \forall \xi^t \in X,$$

$$\Delta^t \in \delta, \quad \Delta'^t \in \delta', \quad \delta' \in LD \setminus \{\delta\}.$$

Доведення: Нехай деяка $\delta' \neq \delta$, $\delta' \in LD_S$ є оптимальною за означенням 1.11, тоді, враховуючи умову маємо.

$$E_y^\delta \frac{1}{1+T} \sum_{t=0}^T r(\xi^t, \Delta^t(\xi^t)) \leq E_y^{\delta'} \frac{1}{1+T} \sum_{t=0}^T r(\xi^t, \Delta'^t(\xi^t)), \quad \forall T \in \mathbb{N}.$$

$$Q_T^\delta(y) \leq Q_T^{\delta'}(y), \quad \forall T \in \mathbb{N}.$$

$$\limsup_{T \rightarrow \infty} (Q_T^\delta(y)) \leq \limsup_{T \rightarrow \infty} (Q_T^{\delta'}(y)).$$

Отже, $p(y) = \limsup_{T \rightarrow \infty} (Q_T^\delta(y))$. Враховуючи припущення $\delta' \neq \delta$ та теорему 1.1 про єдиність оптимальної стратегії з множини LD_S отримаємо протиріччя.

Іншими словами, якщо стратегія δ призводить до найменших однокрокових витрат за довільного стану системи, то δ є оптимальною за визначенням 1.11

Наслідок 2.1 Якщо стратегія δ призводить до найменших однокрокових витрат при будь-якому стані системи, то δ є оптимальною на довільному скінченному горизонті $T \in \mathbb{N}$.

Наслідок 2.2 з доведення твердження 2.1 слідує, що якщо стратегія δ є оптимальною на довільному скінченному горизонті $T \in \mathbb{N}$, то δ є оптимальною за означенням 1.11.

Таким чином, за допомогою твердження 2.1 та наслідків 2.1, 2.2 можливо робити висновки про оптимальність стратегії δ на нескінченному горизонті.

3. РОЗВ'ЯЗАННЯ ЕКОНОМІЧНОЇ ЗАДАЧІ НА СКІНЧЕННОМУ ГОРИЗОНТІ.

Розглянемо числовий приклад задачі сформульованої в роботі [6], та в більш загальному вигляді в роботі [7]. Декілька фірм виробляють два види продукції. Структура графа G визначається, як зображено на рисунку 3.1.

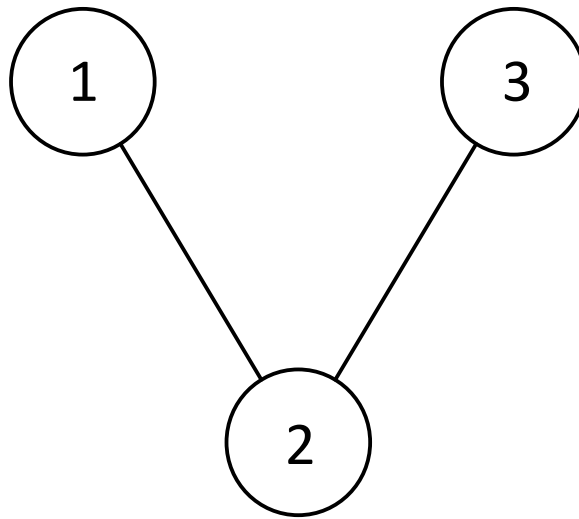


Рисунок 3.1 Структура графа G

Множина вершин: $V = \{1, 2, 3\}$, кожна з яких відображає фірму.

Множина ребер: $B = \{(1, 2), (2, 3)\}$, кожна пара (i, j) , відображає взаємозв'язок між фірмами.

Для кожного елемента $k \in V$ визначено множину станів: $X_k = \{-1, 1\}$, де -1 – виготовляти продукцію 1, 1 – виготовляти продукцію 2.

Для кожної вершини визначено множину можливих станів її повного околу.

$$X_{\tilde{N}(1)} = X_{\tilde{N}(3)} = \{(-1, -1), (-1, 1), (1, -1), (1, 1)\},$$

$$X_{\tilde{N}(2)} =$$

$$= \{(-1, -1, -1), (-1, -1, 1), (-1, 1, -1), (-1, 1, 1), (1, -1, -1), (1, -1, 1), (1, 1, -1), (1, 1, 1)\}.$$

Для кожного елемента $k \in V$ визначено множину допустимих дій $A_k = \{\beta_1, \beta_2\}$, де: β_1 – виконати дію 1, β_2 – виконати дію 2, $\beta_1 = \frac{1}{2}$, $\beta_2 = \frac{1}{4}$.

В роботах [8], [9] локальні ймовірності переходу між станами задані наступним чином.

$$P(y_k | x_{\tilde{N}(k)}) = \frac{e^{-\beta \sum_{j \in N(k)} y_k x_j}}{\sum_{k \in V} e^{-\beta \sum_{j \in \tilde{N}(k)} y_k x_j}};$$

Тоді ймовірності переходу між станами усієї системи.

$$\prod_{k \in V} P(y_k | x_{\tilde{N}(k)}).$$

Де y_k та x_j набувають значень ± 1 , $\beta \in A_k$.

Для кожного елемента $k \in V$ визначено функцію однокрокових винагород в момент часу t .

$$r_k^t(x_k, a_k) = \theta^t \cdot (-\alpha n_{x_k} + \gamma(x_k)).$$

- $\theta = \frac{80}{100}$ – дисконт.
- $\alpha = \frac{1}{4}$ – коефіцієнт втрат через конкуренцію.
- n_{x_k} – кількість фірм, що знаходяться в стані x_k , тобто виробляють таку саму продукцію, $n_{x_k} = 1, 2, \dots, |V|$.
- $\gamma(x_k)$ – дохід від продажу продукції x_k , задано наступним чином.

$$\gamma(x_k) = \begin{cases} \frac{7}{2}, & x_k = 1 \\ \frac{5}{2}, & x_k = -1 \end{cases}.$$

Тоді загальна винагорода в системі.

$$r^t(x, a) = \sum_{i=1}^n \theta^t \cdot (-\alpha n_{x_k} + \gamma(x_k)).$$

Для цієї задачі можливо застосувати підхід, запропонований в [2]. Складність полягає в тому, щоб визначити матрицю переходу Q , оскільки для цього прикладу множини допустимих дій однакові для будь-якого стану системи, тому необхідно визначити $|X \times X \times A| = 512$ значень ймовірності переходу між станами з X за умови прийняття рішення з A . Натомість, для кожної з вершин було обчислено очікувані середні втрати Q_T^δ для скінченного горизонту $T = 10$ за допомогою алгоритму 2.1. Розв'язання задачі в середовищі python наведено в додатку В.

Обчислено однокрокову винагороду в системі $r^t(x, a)$. На останньому кроці $t = T$ для кожного можливого стану x^T . Результат в таблиці 3.1.

Таблиця 3.1 Однокрокова винагорода на останньому кроці T

x^T	(-1,-1,-1)	(-1,-1,1)	(-1,1,-1)	(-1,1,1)	(1,-1,-1)	(1,-1,1)	(1,1,-1)	(1,1,1)
r^T	0.65	0.86	0.86	0.96	0.86	0.97	0.96	0.96

Обчислено очікувані винагороди на кроці $t = T - 1$, для кожного можливого стану системи x^{T-1} . Результат в таблиці 3.2.

Таблиця 3.2 Однокрокова винагорода на кроці $T-1$

x^{T-1}	(-1,-1,-1)	(-1,-1,1)	(-1,1,-1)	(-1,1,1)	(1,-1,-1)	(1,-1,1)	(1,1,-1)	(1,1,1)
$E_{x^{T-1}} r$	1.47	1.95	1.94	2.17	1.95	2.17	2.17	2.15

Обчислено найбільші очікувані винагороди на кожному кроці t в зворотньому порядку, при початковому стані системи y . При цьому, на кожному кроці зафіксовано набір дій $\{a_k\}$, що призводить до найбільшої винагороди. Результат в таблиці 3.3.

Таблиця 3.3 Найбільша очікувана винагорода в момент часу t

$t \backslash y$	(-1, -1, -1)	(-1, -1, 1)	(-1, 1, -1)	(-1, 1, 1)	(1, -1, -1)	(1, -1, 1)	(1, 1, -1)	(1, 1, 1)
10	0.65	0.86	0.86	0.96	0.86	0.97	0.96	0.96
9	1.47	1.95	1.94	2.17	1.95	2.17	2.17	2.15
8	2.49	3.30	3.28	3.67	3.30	3.68	3.67	3.64
7	3.76	4.99	4.97	5.55	4.99	5.56	5.55	5.51
6	5.35	7.10	7.07	7.90	7.10	7.92	7.90	7.84
5	7.34	9.73	9.70	10.83	9.73	10.87	10.83	10.76
4	9.83	13.03	12.99	14.50	13.03	14.55	14.50	14.41
3	12.94	17.16	17.09	19.09	17.16	19.15	19.09	18.97
2	16.82	22.31	22.23	24.83	22.31	24.90	24.83	24.66
1	21.68	28.75	28.65	32.00	28.75	32.09	32.00	31.78
0	27.75	36.80	36.67	40.96	36.80	41.08	40.96	40.69

Оскільки, ймовірності переходу залежать тільки від стану системи та дій, а винагорода залежить від стану системи, то в результаті отримано відповідність між станами системи, та наборами дій, що призводять до найбільших очікуваних винагород:

Таблиця 3.4 Дії a в системі, що призводять до найбільших винагород в стані x

x	(-1, -1, -1)	(-1, -1, 1)	(-1, 1, -1)	(-1, 1, 1)	(1, -1, -1)	(1, -1, 1)	(1, 1, -1)	(1, 1, 1)
a	$\frac{1}{2}, \frac{1}{2}, \frac{1}{2}$	$\frac{1}{2}, \frac{1}{2}, \frac{1}{2}$	$\frac{1}{2}, \frac{1}{2}, \frac{1}{2}$	$\frac{1}{2}, \frac{1}{4}, \frac{1}{4}$	$\frac{1}{2}, \frac{1}{2}, \frac{1}{2}$	$\frac{1}{2}, \frac{1}{4}, \frac{1}{2}$	$\frac{1}{4}, \frac{1}{4}, \frac{1}{2}$	$\frac{1}{4}, \frac{1}{4}, \frac{1}{4}$

Таким чином можливо виділити локальну стратегію δ , що є оптимальною на скінченному горизонті $T = 10$. Стратегія δ визначена в таблиці 3.5

Таблиця 5 Локальна оптимальна стратегія на скінченному горизонті $T=10$

$X_{\tilde{N}(1)}$	δ_1	$X_{\tilde{N}(2)}$	δ_2	$X_{\tilde{N}(3)}$	δ_3
$(-1,-1)$	$\frac{1}{2}$	$(-1,-1,-1)$	$\frac{1}{2}$	$(-1,-1)$	$\frac{1}{2}$
$(-1,1)$	$\frac{1}{2}$	$(-1,-1,1)$	$\frac{1}{2}$	$(-1,1)$	$\frac{1}{2}$
$(1,-1)$	$\frac{1}{2}$	$(-1,1,-1)$	$\frac{1}{2}$	$(1,-1)$	$\frac{1}{2}$
$(1,1)$	$\frac{1}{4}$	$(-1,1,1)$	$\frac{1}{4}$	$(1,1)$	$\frac{1}{4}$
		$(1,-1,-1)$	$\frac{1}{2}$		
		$(1,-1,1)$	$\frac{1}{4}$		
		$(1,1,-1)$	$\frac{1}{4}$		
		$(1,1,1)$	$\frac{1}{4}$		

Враховуючи результати, приведені в таблиці 4 та твердження 2.1, стратегія δ є оптимальною на довільному скінченному горизонті $T \in \mathbb{N}$ та на нескінченному горизонті.

ВИСНОВОК

Результатом проведеної роботи є огляд керованих марковських полів з локальною взаємодією, як інструменту моделювання економічних процесів. Визначено його недоліки, зокрема при моделюванні процесів, для яких множини станів та дій представлені як проміжки на числовій прямій, та переваги – можливість врахувати вплив випадкових процесів на систему та локальну взаємодію елементів.

Запропоновано звуження задачі пошуку оптимальної стратегії до пошуку оптимальної стратегії на скінченному горизонті для процесів, що не потребують планування на нескінченному проміжку часу. Для цього введено визначення оптимальної стратегії на скінченному горизонті, запропоновано алгоритм пошуку оптимальної стратегії на скінченному горизонті, обґрунтовано його дієвість та оцінено складність обчислень. Наведено критерій, за яким можливо визначити, чи буде оптимальна стратегія на скінченному горизонті оптимальною в цілому.

Сформульовано економічну задачу, для вирішення якої застосовано запропонований алгоритм. Побудовано програмну реалізацію алгоритму для розв'язання задачі. Отримані результати проаналізовані на відповідність сформульованому критерію. Подальші дослідження варто присвятити дослідженню існування та єдиності локальних оптимальних стратегій на скінченному горизонті та покращенню алгоритмів пошуку цих стратегій.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

- [1] Chorney R. K., Daduna H., Knopov P. S. Controlled Markov Fields with Finite State Space on Graphs, *Stochastic Models*. 2005. Vol. 21. 1–28p.
- [2] Derman C. Finite State Markovian Decision Processes. *Academic Press*: 1970, 39p.
- [3] Магдич Н. Ю., Керовані Марковські Поля та їх Застосування в Економіці, XIII Всеукраїнська Наукова Конференція Молодих Математиків, Збірка тез С. 32.
- [4] Kauffman S. A. The Origins of Order: Self-Organization and Selection in Evolution. *Oxford University Press*, 1993, 734p.
- [5] Чорней Р. К. Про Деякі Застосування Теорії Керування Марковськими Полями Спеціальної Структури. *УДК 519.217.8*
- [6] David P. A., Foray D. Percolation Structures. Markov Random Fields, The Economics and EDI Standards Diffusion – Center for Economic Policy Research, Stanford University, 1992. 55p.
- [7] Knopov P. S., Some Applied Problems From Random Field Theory, *Cybernetics and Systems Analysis*, 2010, Vol. 46, No. 1.
- [8] Knopov P. S., Markov fields and their applications in economics. *Obchysljuvalna ta Prykladna Matematyka*, 1996, Vol. 80, 33–46p.
- [9] Knopov P. S., and Samosonok A. S., On Markov Stochastic Processes With Local Interactions for Solving Some Applied Problems, *Cybernetics and Systems Analysis*, 2011, Vol. 47, No. 3.

ДОДАТОК А

КЕРОВАНІ МАРКОВСЬКІ ПОЛЯ ТА ЇХ ЗАСТОСУВАННЯ В ЕКОНОМІЦІ

Н.Ю. МАГДИЧ

Керовані поля Маркова є інструментом для побудови моделей, що дозволяють врахувати вплив випадкових процесів на систему та зовнішнє управління, що здійснюється агентами [1]. Таким чином, щоб визначити кероване поле Маркова необхідно визначити Марковське випадкове поле та керування у вершинах.

Означення 1. Марковське випадкове поле – випадкова величина ξ , визначена на $(\Omega, \mathcal{F}, P)$ просторі, що залежить від часу t та задовольняє умову:

$$P(\xi_k^{(t+1)} = y'_k \mid \xi_{(V-k)}^t = x_{(V-k)}^t, \dots) = P(\xi_k^{(t+1)} = y_k \mid \xi_{(N(k))}^t = x_{(N(k))}^t)$$

В роботі [2] доведено існування оптимальної стратегії з множини локальних, стаціонарних, детермінованих стратегій, що є відображеннями множини станів околу вершини на множину допустимих дій агента у вершині.

Означення 2. Марковська локальна, детермінована стратегія – відповідність δ_k між значенням випадкової величини $\xi_{\tilde{N}(k)}^t$ в повному околі вершини k та елементом u_k^t з множини можливих дій в вершині k ;

Модель дозволяє прогнозувати та оптимізувати процеси управління в умовах невизначеності та ризику. Дослідження спирається на теоретичні відомості про контрольовані поля Маркова з локальною, синхронною взаємодією [2]. Ця робота зосереджена на побудові моделей економічних процесів. Головним обмеженням для моделі є складність обчислень, тому розглянута оцінка складності пошуку оптимального локального керування.

ЛІТЕРАТУРА

- [1] Knopov Pavel S. SOME APPLIED PROBLEMS FROM RANDOM FIELD THEORY // Cybernetics and Systems Analysis. – 2010. – Vol. 46, no. 1.
- [2] Ruslan K. Chornei Hans Daduna, Knopov Pavel S. Controlled Markov Fields With Finite State Space on Graphs // Stochastic Models. – 2005. – Vol. 21. – P. 1–28.

Національний університет «Києво-Могилянська академія»
Email address: nazar.mahdych@ukma.edu.ua

ДОДАТОК В

Початкові умови:

```
V=[1,2,3]
B=[(1,2),(2,3)]
```

```
X_1=[-1,1]
X_2=[-1,1]
X_3=[-1,1]
```

```
A_1=[1/2,1/4]
A_2=[1/2,1/4]
A_3=[1/2,1/4]
```

```
X_N_list=[X_N_1, X_N_2, X_N_3]
A_list=[]
```

```
theta=80/100
alpha=1/4
```

Визначення станів повних околів вершин:

```
def x_N_k(k, x):
    if k==0:
        return list((x[0], x[1]))
    if k==1:
        return list((x[0], x[1], x[2]))
    if k==2:
        return list((x[1], x[2]))
```

Визначення функції однокрокових винагород:

```
def gamma(x):
    if x==1:
        return 7/2
    if x==-1:
        return 5/2
```

Однокрокова винагорода для вершини k:

```
def r_k(x_k, a_k, theta, x, t):
    return (theta**t)*(-alpha*(x.count(x_k)) + gamma(x_k))
```

Однокрокова винагорода для всієї системи:

```
def r(x,a, theta, t):
    _r=0
    for i in range(len(x)):
        _r+=r_k(x[i], a[i], theta, x, t)
    return _r
```

Обрахунок ймовірностей переходу між станами системи:

```
def norm(X_i, x_N_i, beta):
    _res=0
    for i in range(len(X_i)):
        _power=0
        for j in range(len(x_N_i)):
            _power+=X_i[i]*x_N_i[j]
        _res+=math.exp(-beta*_power)
    return _res

def local_trans_prob(y_k, x_N_k, X_k, a_k):
    _power=0
    for i in range(len(x_N_k)):
        _power+=y_k*x_N_k[i]
    _numerator=math.exp(-a_k*_power)
    _denominator=norm(X_k, x_N_k, a_k)
    _res=_numerator/_denominator
    return _res

def trans_prob(y, x, a):
    _prob=1
    for i in range(len(V)):
        _prob *= local_trans_prob(y[i], x_N_k(i,x), X_k(i), a[i])
    return _prob
```

Процес пошуку оптимальної стратегії за допомогою алгоритму 2.1:

```
a_t=[]
q_t=[]
_t=[]
T=10

exp_r_next=[0, 0, 0, 0, 0, 0, 0, 0]
for t in range(T,-1,-1):
    a_row=[]
    q_row=[]
```

```

for i in range(len(X)):
    max_exp_r=0
    a_max=[]
    for j in range(len(A)):
        sum_r=0
        exp_r=0
        for k in range(len(X)):
            sum_r += trans_prob(X[k], X[i], A[j])*r(X[k], A[j], theta, t+1)
        exp_r=r(X[i], A[j], theta, t) + sum_r/len(X)+exp_r_next[i]

        if max_exp_r < exp_r:
            max_exp_r=exp_r
            a_max=A[j]
    a_row.append(a_max)
    q_row.append(round(max_exp_r,2))
    exp_r_next[i]=max_exp_r
a_t.append(a_row)
q_t.append(q_row)
_t.append(t)
actions_table = pd.DataFrame(a_t, columns=X, index=_t)
q_table = pd.DataFrame(q_t, columns=X, index=_t)

```