

УДК 681.513

О.О. БРОВAREЦЬ*, О.Ю. САПЕЛЬНИКОВА**

ПРОГРАМНИЙ КОД ТА МЕТОДИКА КЛАСТЕРНОГО АНАЛІЗУ ДАНИХ ІНФОРМАЦІЙНО-ТЕХНІЧНОЇ СИСТЕМИ ТОЧНОГО ЗЕМЛЕРОБСТВА ДЛЯ ВИДІЛЕННЯ ЗОН НЕОДНОРІДНОСТІ АГРОБІОЛОГІЧНОГО СТАНУ БІОРІЗНОМАНІТТЯ СІЛЬСЬКОГОСПОДАРСЬКИХ УГІДЬ

*Київський кооперативний інститут бізнесу і права, м. Київ, Україна

**Київський міжнародний університет, м. Київ, Україна

Анотація. Україна має один із найбільш високих показників родючості у світі. Проте для забезпечення максимальних урожаїв при вирощуванні сільськогосподарських культур важливою складовою є постійний моніторинг агробіологічного стану під час всього періоду вирощування сільськогосподарських культур, до посівів та після збору врожаю. Це дає можливість оптимізувати норми внесення технологічних матеріалів, а відповідно, мінімізувати витрати на вирощування сільськогосподарських культур. Реалізація постійного моніторингу агробіологічного стану сільськогосподарських угідь є важливим елементом сучасного сільськогосподарського виробництва й можлива з використанням широкого спектра пристроїв інформаційно-технічних систем моніторингу агробіологічного стану сільськогосподарських угідь наземного, повітряного та космічного базування на різних етапах вирощування сільськогосподарських культур. У результаті маємо великі масиви даних, які вимагають оперативної їх обробки, візуалізації, представлення для роботи з цими масивами даних та прийняття оперативних рішень для ефективного керування технологічними операціями. Саме тому для роботи з великими масивами даних розроблено методику кластерного аналізу даних інформаційно-технічної системи точного землеробства агробіологічного стану біорізноманіття сільськогосподарських угідь. Запропонований програмний код дозволяє забезпечити диференційоване внесення технологічного матеріалу (насіння, добрив, засобів захисту рослин тощо) на основі даних моніторингу з використанням інформаційно-технічних систем оперативного моніторингу агробіологічного стану сільськогосподарських угідь. За попередніми розрахунками, володіння такою інформацією дозволить на 10–25% зекономити технологічний матеріал і сприятиме підвищенню врожайності сільськогосподарських культур на 10–20 ц/га.

Ключові слова: кластерний аналіз, інформаційно-технічні системи, моніторинг, точне землеробство, біорізноманіття, сільськогосподарські угіддя.

Abstract. Ukraine has one of the highest fertility rates in the world. Therefore, to ensure maximum yields when growing crops, an important component is constant monitoring of the agrobiological state during the entire period of growing crops, before sowing and after harvesting. This makes it possible to optimize the rates of application of technological materials and, accordingly, minimize the costs of growing crops. The implementation of constant monitoring of the agrobiological state of agricultural lands is an important element of modern agricultural production and is possible using a wide range of devices of information and technical systems for monitoring the agrobiological state of agricultural lands, ground, air, and space-based at different stages of growing crops. As a result, we have large data sets that require their prompt processing, visualization, presentation for working with them, and making operational decisions for effective management of technological operations. That is why, to work with large data sets, a method of cluster analysis of data from the information and technical system of precision agriculture of the agrobiological state of the biodiversity of agricultural lands has been developed. The proposed program code allows for the differentiated application of technological material (seeds, fertilizers, plant protection products, etc.) based on monitoring data using information and technical systems for operational monitoring of the agrobiological state of agricultural lands. According to preliminary calculations, possession of such information will allow for saving technological material by 10–25% and help to increase the yield of agricultural crops by 10–20 centners per hectare.

Keywords: cluster analysis, information and technical systems, monitoring, precision agriculture, biodiversity, agricultural lands.

DOI: 10.34121/1028-9763-2025-1-81-90

1. Вступ. Постановка проблеми

Один із фундаментальних процесів у сільськогосподарському виробництві є класифікація даних агробіологічного стану біорізноманіття сільськогосподарських угідь [1–8].

Факти і явища повинні бути впорядковані, класифіковані, перш ніж дослідник зможе зрозуміти їх. Аналіз даних має у своєму розпорядженні великий арсенал методів класифікації.

Залежно від наявності попередньої вибіркової інформації говорять про класифікацію без навчання (інформація про самі класи відсутня) і про класифікацію з навчанням (є навчальні вибірки). Розв’язання задач класифікації без навчання проводиться методами кластерного аналізу.

Мета статті — це розробка програмного коду та методика кластерного аналізу даних інформаційно-технічної системи точного землеробства для виділення зон неоднорідності агробіологічного стану біорізноманіття сільськогосподарських угідь із використанням методу k -середніх для кластерного аналізу.

Постановка задачі — використання методики кластерного аналізу даних, отриманих від інформаційно-технічних систем точного землеробства агробіологічного стану біорізноманіття сільськогосподарських угідь

Позначимо множину агробіологічного стану біорізноманіття сільськогосподарських угідь об’єкта $Q = \{Q_1, Q_2, \dots, Q_N\}$, що складається з N даних, отриманих від інформаційно-технічних систем точного землеробства.

2. Аналіз публікацій за темою дослідження

Дані, отримані від інформаційно-технічних систем точного землеробства, описуються за допомогою n ознак $x_1, x_2, \dots, x_n$. Сукупність значень ознак зведена в матрицю X :

$$X = \begin{bmatrix} x_{11} & x_{12} & \dots & x_{1n} \\ x_{21} & x_{22} & \dots & x_{2n} \\ \dots & \dots & \dots & \dots \\ x_{N1} & x_{N2} & \dots & x_{Nn} \end{bmatrix}.$$

Матрицю X можна інтерпретувати як множину точок $X_1 = (x_{11}, x_{12}, \dots, x_{1n})$, $X_2 = (x_{21}, x_{22}, \dots, x_{2n})$, ..., $X_N = (x_{N1}, x_{N2}, \dots, x_{Nn})$ у n -вимірному евклідовому просторі E_n .

Нехай m — ціле число, $m < N$. Тоді задачу кластерного аналізу можна сформулювати так: на підставі даних, що містяться в x , розбити множину об’єктів o на m кластерів (підмножин) $K_1, K_2, \dots, K_m$ так, щоб кожний об’єкт належав одній і тільки одній підмножині, і щоб об’єкти, що належать тому самому кластеру, були подібними (близькими), тоді як об’єкти, що належать різним кластерам, були несхожими (відмінними).

Розв’язання задачі кластерного аналізу вимагає формалізації понять близькості й відмінності. Розглянемо поняття відстані між точками X_i і X_j .

3. Функції відстані й подібності

Невід’ємна дійсна функція $d(X_i, X_j)$ називається функцією відстані (метрикою), якщо виконуються такі властивості:

1. $d(X_i, X_j) \geq 0 \forall X_i, X_j \in E_n$.
2. $d(X_i, X_j) = 0$ лише $X_i = X_j$.

$$3. d(X_i, X_j) \leq d(X_j, X_i).$$

4. $d(X_i, X_j) \leq d(X_i, X_k) + d(X_k, X_j)$, де X_i, X_j, X_k — будь-які три точки з E_n (так зване «правило трикутника»).

Значення функції d для двох точок X_i, X_j еквівалентно відстані між об'єктами O_i і O_j .

Як приклад функцій відстані приведемо найбільш уживані.

Евклідова відстань:

$$d(X_i, X_j) = \sqrt{\sum_{k=1}^n (x_{ik} - x_{jk})^2}.$$

Сума абсолютних відхилень:

$$d(X_i, X_j) = \sum_{k=1}^n |x_{ik} - x_{jk}|.$$

Відстані між об'єктами можуть бути зведені у квадратну симетричну матрицю відстаней:

$$D = \begin{bmatrix} 0 & d_{12} & \cdots & d_{1N} \\ d_{21} & 0 & \cdots & d_{2N} \\ \cdots & \cdots & \cdots & \cdots \\ d_{N1} & d_{N2} & \cdots & 0 \end{bmatrix}.$$

Поняттями, протилежними відстані, є поняття подібності. Мірою подібності називають невід'ємну дійсну функцію s , що задовольняє таким аксіомам:

$$1. 0 \leq s(X_i, X_j) \leq 1, X_i = X_j.$$

$$2. s(X_i, X_j) = 1.$$

$$3. s(X_i, X_j) = s(X_j, X_i).$$

Значення функцій подібності елементів множини O можна об'єднати в матрицю подібності:

$$S = \begin{bmatrix} 1 & s_{12} & \cdots & s_{1N} \\ s_{21} & 1 & \cdots & s_{2N} \\ \cdots & \cdots & \cdots & \cdots \\ s_{N1} & s_{N2} & \cdots & 1 \end{bmatrix}.$$

4. Міра близькості між кластерами

Нехай K_i — i -й кластер, що містить N_i об'єктів;

$\bar{X}_i$ — середнє арифметичне спостережень, що входять у K_i , тобто

$$\bar{X}_i = \frac{1}{N_i} \sum_{j=1}^{N_i} X_j, X_j \in K_i.$$

5. Найбільш уживані відстані між кластерами

Відстань, вимірювана за принципом найближчого сусіда:

$$D_{\min}(K_l, K_m) = \min_{X_i \in K_l, X_j \in K_m} d(X_l, X_m).$$

Відстань, вимірювана за принципом далекого сусіда:

$$D_{\max}(K_l, K_m) = \max_{X_i \in K_l, X_j \in K_m} d(X_l, X_m).$$

Відстань, вимірювана по центрах ваги кластерів:

$$D_{\text{cp}}(K_l, K_m) = d(\bar{X}_l, \bar{X}_m).$$

Ілюстрація трьох наведених мір:

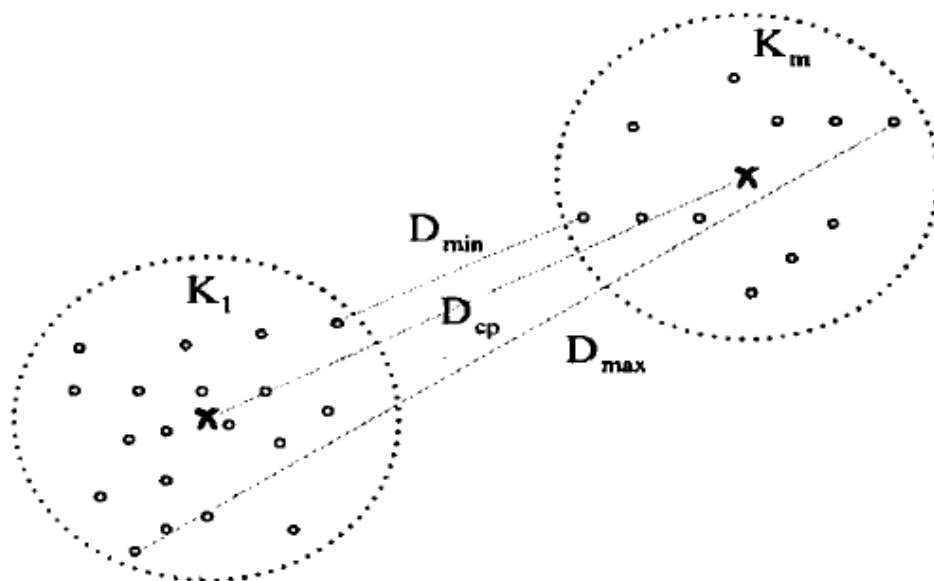


Рисунок 1 — Різні міри близькості між кластерами

6. Функціонали якості розбиття на кластери

Розбиття вихідної сукупності об'єктів на кластери може здійснюватися різними способами. Природно визначити той кількісний критерій, виходячи з якого можна було б робити висновок про якість розбиття. З цією метою в задачах кластерного аналізу вводять поняття функціонала якості розбиття $G(K)$. Цей функціонал визначається на множині всіх можливих варіантів розбиття. Розбиття k^* , що приводить до екстремуму обраного функціонала, вважається найкращим.

Слід зазначити, що вибір функціонала, так само як і вибір метрики, здійснюється довільно.

Функціонал якості розбиття повинен відбивати різноманітність множини об'єктів. Користуючись поняттям відстані, введемо спочатку міри розсіювання множини об'єктів, які будуються на базі матриці відстаней D . Величину $S_d(O)$, рівну напівсумі всіх елементів матриці D , називають загальним розсіюванням, а величину

$$\bar{S}_d(O) = \frac{S_d(O)}{N_d},$$

де

$$N_d = \frac{N(n-1)}{2},$$

називають середнім розсіюванням множини O .

Визначимо функціонал якості розбиття. Нехай $K = \{K_1, K_2, \dots, K_m\}$ — деяке фіксоване розбиття спостережень $X_1, X_2, \dots, X_N$ на задане число m кластерів $K_1, K_2, \dots, K_m$. Як функціонали часто обираються такі характеристики.

Сума загальних розсіювань:

$$G_1(K) = \sum_{i=1}^m S_d(K_i).$$

Сума середніх розсіювань:

$$G_2(K) = \sum_{i=1}^m \bar{S}_d(K_i).$$

7. Алгоритми роздільної кластеризації

Задача кластерного аналізу носить комбінаторний характер. Прямий спосіб розв'язання такої задачі полягає в повному переборі всіх можливих варіантів розбиття на кластери й виборі розбивки, що забезпечує екстремальне значення функціонала. Такий спосіб розв'язання називають кластеризацією повним перебором. Аналогом кластерної проблеми комбінаторної математики є задача розбиття множини з n об'єктів на m підмножин. Число таких варіантів розбиття позначається через $S(n, m)$ і називається числом Стірлінга другого роду. Кластеризація повним перебором можлива в тих випадках, коли число об'єктів і кластерів не дуже велике.

Найбільш широке застосування одержали алгоритми послідовної кластеризації. В цих алгоритмах проводиться послідовний вибір чергових спостережень і для кожного з них вирішується питання, до якого з m кластерів його віднести. Розглянемо одну з послідовних кластер-процедур — метод k -середніх.

Метод k -середніх

Особливість цього алгоритму полягає в тому, що він виконується в два етапи: на першому етапі в просторі E_n виконується пошук точок — центрів кластерів, а потім спостереження розподіляються по тих кластерах, до центрів яких вони тяжіють. Алгоритм працює в припущенні, що число m кластерів відомо. Перший етап починається з відбору m об'єктів, які приймаються як нульове наближення центрів кластеризації. Це можуть бути перші m зі списку об'єктів або випадково відібрані.

Кожному центру приписується одинична вага. На першому кроці алгоритму отримується перша з точок, що залишилися (позначимо її X_{m+1}), і з'ясовується, до якого з центрів вона виявилася ближче всього в сенсі обраної d метрики. Цей центр замінюється новим, обумовленим як зважена комбінація старого центра й нової точки. Вага центра збільшується на одиницю і т.д., доки не будуть пройдені всі точки. Метод k -середніх є найбільш точним для проведення кластерного аналізу, то його і будемо використовувати для кластерного аналізу даних інформаційно-технічної системи точного землеробства при виділенні зон неоднорідності агробіологічного стану біорізноманіття сільськогосподарських угідь.

Ієрархічний кластерний аналіз

Поряд із роздільним кластерним аналізом широко застосовується ієрархічний кластерний аналіз, мета якого полягає в одержанні всієї ієрархії об'єктів, а не окремого розбиття. Якщо використати неорієнтований граф, то структура такої розбивки являє собою дерево. Сам процес класифікації представляє собою побудову ієрархічного дерева досліджуваної сукупності об'єктів. Графічне зображення неорієнтованого графа ієрархії на площині називають дендрограмою.

8. Кластерний аналіз в R

Основні процедури кластерного аналізу реалізовані в R у пакеті `stats, cluster`. Є також і інші пакети, які необхідно завантажити і встановити окремо. Наприклад, `pvcust`.

9. Виклад основного змісту дослідження

Сучасне землеробство передбачає широке використання ефективних технологій точного землеробства, які дають можливість раціонально використовувати технологічний матеріал з використанням технологій змінних норм їх внесення. Проте реалізація технології змінних норм внесення технологічних матеріалів передбачає володіння великим масивом даних, отриманих від систем моніторингу агробіологічного стану сільськогосподарських угідь. Тому для ефективного використання даної технології виникає необхідність у розробці

методики кластерного аналізу даних інформаційно-технічної системи точного землеробства для виділення зон неоднорідності агробіологічного стану сільськогосподарських угідь, що є основою змінних норм внесення технологічного матеріалу. Така методика дозволяє представити дані, візуалізувати їх та використати для технологій змінних норм внесення технологічного матеріалу [9–12].

При моніторингу агробіологічного стану сільськогосподарських угідь отримуємо масив даних із параметрами агробіологічного стану сільськогосподарських угідь із прив'язкою до географічних координат (довготи і широти). Проте прийняття їх неможливе без подальшої математичної обробки та візуалізації даних для подальшого прийняття рішень щодо ефективного керування агробіологічним станом сільськогосподарських угідь.

З цією метою необхідно написати програмний код із використанням методики кластерного аналізу даних інформаційно-технічної системи точного землеробства для виділення зон неоднорідності агробіологічного стану біорізноманіття сільськогосподарських угідь. Як агробіологічні дані будемо використовувати такі дані: рН(кислотність), вологість, родючість, N, P, K, електропровідність.

Використовуючи мову програмування R для кластерного аналізу агробіологічних даних (визначення рН, вологість, родючість, N, P, K, електропровідність) із прив'язкою до довжини та ширини для визначення зони неоднорідності сільськогосподарських угідь, включаємо реальний приклад із гіпотетичними географічними даними, які будуть імітувати польові виміри параметрів агробіологічного стану. Використовуємо дані для створення карт зон неоднорідності з використанням інтерполяції. Напишемо програму на мові R для кластерного аналізу агробіологічних даних для визначення зони неоднорідності сільськогосподарських угідь на основі даних з Excel у форматі csv.

Програма на мові R для кластерного аналізу агробіологічних даних із прив'язкою до довготи та широти, яка враховує значення рН, вологість, родючість, N, P, K та електропровідність.

Технічне завдання до розробки програми на мові R:

1. Завантаження даних із CSV.
2. Попередня обробка даних.
3. Кластеризація агробіологічних параметрів.
4. Інтерполяція для створення зони неоднорідності.
5. Генерація карти.

Код на R

```
# Встановлення необхідних пакетів
install.packages(c("tidyverse", "sf", "factoextra", "cluster", "gstat", "raster", "readr"))
library(tidyverse)
library(sf)
library(factoextra)
library(cluster)
library(gstat)
library(raster)
library(readr)

# 1. Завантаження даних із CSV
# Імпортуємо дані (замість "data.csv" вкажіть свій файл)
file_path
<- "data.csv" # Шлях до вашого файлу
data <- read_csv(file_path)
```

```

# Попередній огляд даних
glimpse(data)

# Очікується, що CSV має стовпці: longitude, latitude, pH, moisture, fertility, N, P, K,
conductivity

# 2. Масштабування агробіологічних параметрів
scaled_data <- scale(data %>% select(pH, moisture, fertility, N, P, K, conductivity))

# 3. Кластерний аналіз
# Визначення оптимальної кількості кластерів методом "Elbow"
fviz_nbclust
fviz_ (scaled_data, kmeans, method = "wss") +
  labs(title = "Оптимальна кількість кластерів")

# Використання k-means для 3 кластерів (змінити кількість за потреби)
set.seed(123)
kmeans_result <- kmeans(scaled_data, centers = 3, nstart = 25)

# Додавання кластерів до даних
data$cluster <- as.factor(kmeans_result$cluster)

# 4. Геопросторова інтерполяція
# Перетворення даних у просторовий формат sf
geo_data <- st_as_sf(data, coords = c("longitude", "latitude"), crs = 4326)

# Створення сітки для інтерполяції
grid <- st_make_grid(geo_data, cellsize = 0.01, what = "centers")
grid_sf <- st_as_sf(grid)

# Інтерполяція значень кластерів методом IDW
idw_result <- gstat::idw(formula = cluster ~ 1, locations = geo_data, newdata = grid_sf)

# Перетворення результатів інтерполяції у растровий формат
raster_result <- raster::rasterFromXYZ(as.data.frame(cbind(st_coordinates(grid_sf),
idw_result$var1.pred)))

# 5. Візуалізація результатів
# Карта зон неоднорідності
ggplot() +
  geom_raster(data = as.data.frame(raster::rasterToPoints(raster_result)),
    aes(x = x, y = y, fill = layer)) +
  scale_fill_viridis_d

(name = "Кластер") +
  geom_point(data = data, aes(x = longitude, y = latitude, color = cluster), size = 3) +
  labs(title = "Карта зон неоднорідності сільськогосподарських угідь",
    x = "Довгота", y = "Широта", color = "Кластер") +
  theme_minimal()

```

Опис роботи програми кластерного аналізу даних інформаційно-технічної системи точного землеробства для виділення зон неоднорідності агробіологічного стану біорізноманіття сільськогосподарських угідь.

1. Завантаження даних: програма працює з файлом у форматі CSV, який містить координати (довгота, широта) і агробіологічні параметри: pH, `moisture` (вологість), N (азот), P (фосфор), K (калій), EMc (електропровідність).

2. Масштабування: масштабуються всі класи.

3. Кластеризація: *k*-means, *a* оптимальна кількість класів.

4. Інтерполяція: зони кластерів IDW (інверсна вага відстані), щоб побудувати кластерний аналіз.

5. Візуалізація: карта класів кластерного аналізу.

Для реалізації поставленої мети реалізовано програмний код у мові програмування R для визначення зони неоднорідності на основі геовизначених параметрів (детермінованих) — даних агробіологічного стану у визначених місцях та показати їх.

Для аналізу зон неоднорідності на основі геоприв'язаних даних параметрів агробіологічного стану можна використовувати мову R із такими бібліотеками, як *sf* для роботи з геопросторовими даними, *ggplot2* для візуалізації, *spatstat* для аналізу просторових даних або *gstat* для інтерполяції.

Програмний код у мові програмування R для аналізу та візуалізації зон неоднорідності.

```
# Встановлення необхідних бібліотек
install.packages(c("sf", "ggplot2", "sp", "gstat", "raster"))

# Завантаження бібліотек
library(sf)          # Робота з геопросторовими даними
library(ggplot2)     # Візуалізація
library(gstat)       # Інтерполяція
library(raster)      # Робота з растровими даними

# Дані (створення штучного набору для прикладу)
# Створимо набір даних із координатами (широта і довгота) та агробіологічним показником
set.seed(123)
data <- data.frame(
  x = runif(100, 30, 40),  # Довгота
  y = runif(100, 50, 60),  # Широта
  agro_param = runif(100, 10, 50) # Параметр агробіологічного стану
)

# Перетворення даних у геопросторовий об'єкт
geo_data <- st_as_sf(data, coords = c("x", "y"), crs = 4326)

# Інтерполяція значень для створення зон неоднорідності
# Перетворення в проекцію для аналізу
geo_data_proj <- st_transform(geo_data, crs = 32636) # UTM Zone 36N

# Створення сітки для інтерполяції
grd <- st_make_grid(geo_data_proj, cellsize = 500, what = "centers")
grd_sf <- st_as_sf(grd)

# Інтерполяція методом Крігінга
kriging_model <- gstat(formula = agro_param ~ 1, data = geo_data_proj)
```



```
kriging_result <- predict(kriging_model, newdata = grd_sf)

# Перетворення результату у растровий формат
rast <- rasterFromXYZ(as.data.frame(st_coordinates(grd_sf), kriging_result$var1.pred))

# Візуалізація зон неоднорідності
ggplot() +
  geom_raster(data = as.data.frame(rasterToPoints(rast)),
    aes(x = x, y = y, fill = layer)) +
  geom_point(data = data, aes(x = x, y = y), color = "red", size = 1) +
  scale_fill_viridis_c(name = "Agro param") +
  labs(title = "Зони неоднорідності агробіологічного стану") +
  theme_minimal()
```

Пояснення коду

Генерація даних: створюються координати і показники.

1. Перетворення в геопросторовий формат: використовується `sf` для роботи з геодами.
2. Інтерполяція: застосовується метод Кригінга для просторового згладжування даних.
3. Візуалізація: використовується `ggplot2` для відображення зон неоднорідності на карті.

10. Висновок

Код створює карту, яка відображає зони неоднорідності агробіологічного стану на основі даних.

Написана програма на мові R для визначення зони неоднорідності на основі геодетермінованих даних із параметрами агробіологічного стану у визначених місцях та показати її реалізацію на конкретному прикладі.

Зазначена програма містить реальний приклад із гіпотетичними геоданими, які будуть імітувати польові виміри параметрів агробіологічного стану.

Використано дані для створення карти зон неоднорідності з використанням інтерполяції.

СПИСОК ДЖЕРЕЛ

1. Федірець О.В., Броварець О.О., Мартин О.М. Інноваційне управління агропродовольчою сферою як соціально-економічною системою в умовах безпекових викликів. *Інвестиції: практика та досвід. Економічна наука*. 2024. № 8. С. 52–58.
2. Броварець О.О. Аналіз зон неоднорідності (варіабельності) стану сільськогосподарських угідь із використанням засобів моніторингу та розробка програмного забезпечення для його забезпечення. *Математичні машини і системи*. 2022. № 3. С. 85–99.
3. Броварець О. Розумні машини для розумних господарів. *Зерно*. 2016. № 9 (81). С. 262–266.
4. Броварець О. Від безплужного до глобального розумного землеробства. *Техніка і технології АПК*. 2016. № 9 (84). С. 19–23.
5. Броварець О. Від безплужного до глобального розумного землеробства. *Техніка і технології АПК*. 2016. № 10 (85). С. 28–30.
6. Адамчук В.В., Мойсеєнко В.К., Кравчук В.І., Войтюк Д.Г. Техніка для землеробства майбутнього. *Механізація та електрифікація сільського господарства*. Глеваха: ННЦ «ІМЕСГ», 2002. Вип. 86. С. 20–32.
7. Сучасні тенденції розвитку конструкцій сільськогосподарської техніки / За ред. В.І. Кравчука, М.І. Грицишина, С.М. Ковалю. К.: Аграрна наука, 2004. 398 с.

8. Мироненко В.Г., Броварець О.О. Інтегровані системи автоматичного управління технологічними процесами у рослинництві. *Науковий вісник Національного університету біоресурсів і природокористування України*. К., 2015. Ч. 1, Вип. 212. С. 62–71.
9. Ковалюк Т.В. Основи програмування. К.: ВНЗ, 2005. 384 с.
10. Яшина О.В., Жульковський О.О. Обчислювальна техніка та програмування. Дніпродзержинськ: ДДТУ, 2007. 309 с.
11. Лавріщева К.М. Основні напрями досліджень у програмній інженерії і шляхи їхнього розвитку. *Проблеми програмування*. 2003. № 3–4. С. 44–58.
12. Бабенко Л.П., Лавріщева К.М. Основи програмної інженерії: навч. посібн. К.: Знання, 2001. 269 с.

Стаття надійшла до редакції 31.12.2024