

Міністерство освіти і науки України
Тернопільський національний технічний університет
імені Івана Пулюя

Кафедра вищої математики

Методичні вказівки

до виконання практичних робіт з дисципліни

«Статистичні методи обробки результатів досліджень»

для здобувачів вищої освіти другого (магістерського) рівня
за спеціальністю G13 «Харчові технології»
денної та заочної форм навчання

Тернопіль, 2026

Габрусев Г. В. Методичні вказівки до виконання практичних робіт з дисципліни «Статистичні методи обробки результатів досліджень» для здобувачів вищої освіти другого (магістерського) рівня за спеціальністю G13 «Харчові технології» денної та заочної форм навчання / Григорій Габрусев. — Тернопіль : СМП «Тайп», 2026. — 57 с.

Укладач:

Г. В. Габрусев, кандидат фізико-математичних наук, доцент.

Рецензент:

Б. Г. Шелестовський, кандидат фізико-математичних наук, доцент.

Методичні вказівки розглянуті і затверджені на засіданні
кафедри вищої математики ТНТУ ім. Івана Пулюя
(Протокол № 1 від 27 серпня 2026 року).

Схвалені і рекомендовані до друку методичною комісією
факультету прикладних інформаційних технологій та електроінженерії
Тернопільського національного технічного університету імені Івана Пулюя
(Протокол № 1 від 28 серпня 2026 року).

© Габрусев Г. В.
© СМП "ТАЙП", 2026

Передмова

Практикум призначений для послідовного опанування статистичних методів, що застосовуються під час обробки результатів експериментальних і виробничих досліджень у харчових технологіях. Матеріал побудовано від імовірнісних моделей і первинного аналізу вибірок до перевірки статистичних гіпотез, ANOVA, кореляційного та регресійного аналізу.

Кожна практична робота поєднує математичний зміст методу, ручне виконання ключових етапів та комп'ютерну перевірку. Основним середовищем розрахунків є Microsoft Excel. Для процедур, які недоцільно розгортати вручну, використовуються Real Statistics Resource Pack або Jamovi.

постановка задачі → вибір методу → ручний розрахунок → Excel → зіставлення → висновок

Головний принцип. Правильне числове значення саме по собі не є завершеним статистичним результатом. Студент повинен обґрунтувати вибір методу, перевірити умови його застосування та пояснити предметний зміст результату.

Структура практичних робіт

№	Назва	Основний результат
1	Елементи теорії ймовірностей та комбінаторики в задачах контролю якості	Вибір імовірнісної моделі та розрахунок ймовірностей
2	Робота з вибірками, побудова гістограм та контрольних карт Шухарта	Первинний опис даних і статистична стабільність процесу
3	Розрахунок описової статистики та довірчих інтервалів	Точкове та інтервальне оцінювання параметрів
4	Комплексний порівняльний аналіз двох технологічних процесів	Нормальність, дисперсії, Student/Welch, CI, Cohen's d
5	Органолептичний аналіз непараметричними критеріями	Mann–Whitney і Wilcoxon
6	Дисперсійний (ANOVA) та кореляційний аналіз технологічних даних	ANOVA як основна частина; кореляція як коротке продовження
7	Побудова та оцінювання регресійної моделі	Лінійна модель, залишки, значущість і прогноз

Єдина структура кожної практичної роботи

1. **1. Мета роботи.** Конкретні результати навчання.
2. **2. Короткі теоретичні відомості.** Означення, формули, умови застосування та алгоритм вибору методу.
3. **3. Порядок виконання роботи.** Послідовний алгоритм від вихідних даних до висновку.
4. **4. Приклад виконання — демонстраційний варіант 0.** Повний покроковий розв'язок.
5. **5. Виконання та перевірка в Microsoft Excel.** Комп'ютерний етап після ручного розрахунку.
6. **6. Вимоги до оформлення звіту.** Уніфіковані для серії.
7. **7. Контрольні питання.** Поняття, умови застосування та інтерпретація.

Загальні вимоги до висновків

Статистичний висновок повинен містити рішення щодо гіпотези або моделі, основну статистику, p-value, довірчий інтервал та/або величину ефекту, якщо вони передбачені роботою.

Предметний технологічний висновок повинен пояснювати, що встановлений статистичний результат означає для продукту, рецептури, режиму, показника якості або технологічного процесу.

Зміст

Практична робота №1. Елементи теорії ймовірностей та комбінаторики в задачах контролю якості	4 ст.
Практична робота №2. Робота з вибірками, побудова гістограм та контрольних карт Шухарта	10 ст.
Практична робота №3. Розрахунок описової статистики та довірчих інтервалів	17 ст.
Практична робота №4. Комплексний порівняльний аналіз двох технологічних процесів	24 ст.
Практична робота №5. Органолептичний аналіз непараметричними критеріями	31 ст.
Практична робота №6. Дисперсійний (ANOVA) та кореляційний аналіз технологічних даних	39 ст.
Практична робота №7. Побудова та оцінювання регресійної моделі	47 ст.
Рекомендована література.	57 ст.

Практична робота №1

Елементи теорії ймовірностей та комбінаторики в задачах контролю якості

1. Мета роботи

Закріпити основні поняття комбінаторики та теорії ймовірностей і набути практичних навичок їх застосування до задач вибіркового контролю якості та аналізу випадкових подій у харчових технологіях.

- розпізнавати комбінаторну схему задачі та обчислювати число можливих вибірок, розміщень і комбінацій;
- обчислювати ймовірності випадкових подій за класичним означенням, теоремами додавання і множення;
- застосовувати формулу повної ймовірності та формулу Байєса для аналізу джерел технологічного браку;
- використовувати схему Бернуллі для оцінювання ймовірності появи певної кількості дефектних одиниць продукції;
- обґрунтовувати вибір ймовірнісної моделі та інтерпретувати результат у контексті контролю якості.

2. Короткі теоретичні відомості

2.1. Основні поняття теорії ймовірностей

Випадковим експериментом називають експеримент, результат якого неможливо однозначно передбачити до його проведення. Прикладами є вибір одиниці продукції з партії, визначення відповідності упаковки нормативній масі, реєстрація наявності дефекту, лабораторний аналіз випадково відібраної проби.

Окремий можливий результат випадкового експерименту називають **елементарною подією**. Множину всіх можливих елементарних результатів називають **простором елементарних подій**. Подія може бути достовірною, неможливою або випадковою.

2.2. Основи комбінаторики

Комбінаторні формули використовують для визначення кількості можливих способів вибору або впорядкування об'єктів.

Факторіал:

$$n! = 1 \cdot 2 \cdot 3 \cdots n, \quad 0! = 1.$$

Перестановки — коли використовуються всі n різних об'єктів і порядок важливий:

$$P_n = n!.$$

Розміщення — коли з n об'єктів вибирають k і порядок має значення:

$$A_n^k = \frac{n!}{(n-k)!}.$$

Комбінації — коли з n об'єктів вибирають k і порядок не має значення:

$$C_n^k = \frac{n!}{k!(n-k)!}.$$

Умова задачі	Формула
Використовуються всі об'єкти, порядок важливий	P_n
Вибирають k із n , порядок важливий	A_n^k
Вибирають k із n , порядок неважливий	C_n^k

2.3. Класичне означення ймовірності

Якщо всі елементарні результати випадкового експерименту є рівноможливими, ймовірність події A визначають за формулою:

$$P(A) = \frac{m}{n},$$

де n — загальна кількість рівноможливих результатів, m — кількість результатів, сприятливих події A . Обов'язково:

$$0 \leq P(A) \leq 1, \quad P(\Omega) = 1, \quad P(\emptyset) = 0.$$

2.4. Протилежна подія

Подію, яка відбувається тоді й тільки тоді, коли подія A не відбулася, називають протилежною і позначають $\bar{A}$.

$$P(\bar{A}) = 1 - P(A).$$

Ця формула особливо зручна для задач типу «знайти ймовірність появи хоча б одного дефектного виробу»:

$$P(X \geq 1) = 1 - P(X = 0).$$

2.5. Теорема додавання ймовірностей

Для двох довільних подій:

$$P(A \cup B) = P(A) + P(B) - P(A \cap B).$$

Якщо події несумісні, $P(A \cap B) = 0$, тому:

$$P(A \cup B) = P(A) + P(B).$$

2.6. Умовна ймовірність

$$P(A | B) = \frac{P(A \cap B)}{P(B)}, \quad P(B) > 0.$$

У технологічному аналізі умовна ймовірність може відповідати запитанню: «Яка ймовірність того, що продукт виготовлено на певній лінії, якщо встановлено, що він має дефект?»

2.7. Теорема множення

$$P(A \cap B) = P(A)P(B | A).$$

Якщо події незалежні:

$$P(A \cap B) = P(A)P(B).$$

Останню формулу не можна механічно застосовувати до залежних подій, наприклад до послідовного відбору з малої партії без повернення.

2.8. Вибірковий контроль без повернення

Нехай партія містить N одиниць продукції, серед яких D дефектних і $N - D$ придатних. Без повернення відбирають n одиниць. Ймовірність того, що серед них буде рівно k дефектних:

$$P(X = k) = \frac{C_D^k C_{N-D}^{n-k}}{C_N^n}.$$

Чисельник визначає число вибірок потрібного складу, знаменник — загальне число можливих вибірок обсягу n .

2.9. Формула повної ймовірності

Нехай подія A може відбутися внаслідок реалізації однієї з взаємовиключних гіпотез $H_1, H_2, \dots, H_m$. Тоді:

$$P(A) = \sum_{i=1}^m P(H_i)P(A | H_i).$$

Наприклад, якщо продукція надходить з кількох ліній із різними рівнями дефектності, формула дозволяє знайти загальну ймовірність отримання дефектної одиниці.

2.10. Формула Байєса

Якщо подія A уже відбулася, формула Байєса дозволяє переоцінити ймовірність відповідної гіпотези:

$$P(H_i | A) = \frac{P(H_i)P(A | H_i)}{P(A)}.$$

У контролі якості це дає змогу оцінити найбільш імовірне джерело виявленого браку.

2.11. Повторні незалежні випробування. Схема Бернуллі

Нехай проводиться n незалежних випробувань. У кожному подія A відбувається з однаковою ймовірністю p , а не відбувається з ймовірністю $q = 1 - p$. Тоді ймовірність рівно k появ:

$$P_n(k) = C_n^k p^k q^{n-k}.$$

Жодної появи:

$$P_n(0) = q^n = (1 - p)^n.$$

Хоча б одна поява:

$$P(X \geq 1) = 1 - (1 - p)^n.$$

Необхідні умови схеми Бернуллі: фіксована кількість випробувань, два результати відносно події, незалежність випробувань, сталість p .

2.12. Алгоритм вибору ймовірнісної моделі

- Задача на кількість способів $\rightarrow$ визначити, чи важливий порядок $\rightarrow P_n, A_n^k$ або C_n^k .
- Вибір із кінцевої партії без повернення $\rightarrow$ комбінаторна модель вибірки.
- Серія незалежних однотипних випробувань зі сталим $p \rightarrow$ схема Бернуллі.
- Подія може виникнути з декількох джерел $\rightarrow$ формула повної ймовірності.
- Результат уже відомий і потрібно оцінити його джерело $\rightarrow$ формула Байєса.

3. Порядок виконання роботи

1. Проаналізувати постановку кожної задачі та визначити випадкову подію.
2. Для комбінаторної задачі встановити загальну кількість об'єктів, кількість вибраних об'єктів і те, чи має значення порядок.
3. Для вибіркового контролю встановити N, D, n, k та спосіб відбору — з поверненням або без.
4. Обчислити ймовірність події вручну.
5. Для задачі з кількома джерелами сформулювати гіпотези H_i , визначити апіорні та умовні ймовірності, застосувати повну ймовірність і, за необхідності, Байєса.
6. Для серії повторних випробувань перевірити умови схеми Бернуллі.
7. Повторити основні обчислення в Microsoft Excel.
8. Зіставити результати ручного та комп'ютерного розрахунків.
9. Для кожної задачі сформулювати короткий предметний висновок.

4. Приклад виконання. Демонстраційний варіант 0

Завдання 1. Формування вибірки для лабораторного контролю

З партії відібрано 10 контрольних зразків продукту. Для поглибленого лабораторного аналізу необхідно вибрати 3 з них.

1. Скільки різних груп із трьох зразків можна сформувати, якщо порядок їх вибору не має значення?

Оскільки порядок не має значення, використовуємо комбінації:

$$C_{10}^3 = \frac{10!}{3!(10-3)!} = \frac{10!}{3!7!}.$$

Скорочуємо факторіали:

$$C_{10}^3 = \frac{10 \cdot 9 \cdot 8}{3 \cdot 2 \cdot 1} = 120.$$

Отже, можна сформувати **120 різних груп**.

2. Скількома способами можна послідовно вибрати три зразки для трьох різних лабораторних операцій?

Тепер порядок важливий, тому використовуємо розміщення:

$$A_{10}^3 = \frac{10!}{7!} = 10 \cdot 9 \cdot 8 = 720.$$

Отже, існує **720 способів**.

Висновок. Для звичайного формування вибірки порядок неважливий і застосовують комбінації; якщо вибрані об'єкти призначаються до різних ролей/операцій, порядок стає важливим.

Завдання 2. Контроль партії без повернення

У партії міститься 20 упаковок продукції, серед яких 4 мають дефект маркування. Для контролю випадково без повернення відбирають 3 упаковки.

Маємо: $N = 20$, $D = 4$, $N - D = 16$, $n = 3$.

Загальна кількість способів вибрати три упаковки:

$$C_{20}^3 = \frac{20!}{3!17!} = 1140.$$

1. Ймовірність рівно однієї дефектної упаковки.

Одну дефектну упаковку можна вибрати C_4^1 способами, дві придатні — C_{16}^2 способами. Тому:

$$P(X = 1) = \frac{C_4^1 C_{16}^2}{C_{20}^3}.$$

$$C_4^1 = 4, \quad C_{16}^2 = \frac{16 \cdot 15}{2} = 120.$$

$$P(X = 1) = \frac{4 \cdot 120}{1140} = \frac{480}{1140} \approx 0.4211.$$

Ймовірність становить приблизно **42.11%**.

2. Ймовірність хоча б однієї дефектної упаковки.

Зручніше використати протилежну подію. Спочатку визначимо ймовірність того, що всі три упаковки придатні:

$$P(X = 0) = \frac{C_{16}^3}{C_{20}^3} = \frac{560}{1140} \approx 0.4912.$$

Тоді:

$$P(X \geq 1) = 1 - P(X = 0) = 1 - 0.4912 = 0.5088.$$

Предметний висновок. Навіть за наявності 20% дефектних упаковок у партії вибірка лише з трьох одиниць виявляє хоча б один дефект приблизно у половині випадків. Ефективність вибіркового контролю суттєво залежить від обсягу вибірки.

Завдання 3. Визначення ймовірного джерела браку

Продукція надходить із двох фасувальних ліній. Лінія №1 виготовляє 60% продукції, лінія №2 — 40%. Частка дефектної продукції: 2% для лінії №1 і 5% для лінії №2.

Позначимо H_1 — продукт із лінії 1, H_2 — продукт із лінії 2, D — продукт дефектний.

$$P(H_1) = 0.60, \quad P(H_2) = 0.40, \quad P(D | H_1) = 0.02, \quad P(D | H_2) = 0.05.$$

1. Загальна ймовірність дефекту.

$$P(D) = P(H_1)P(D | H_1) + P(H_2)P(D | H_2).$$

$$P(D) = 0.60 \cdot 0.02 + 0.40 \cdot 0.05 = 0.012 + 0.020 = 0.032.$$

Отже, загальний рівень дефектності становить 3.2%.

2. Якщо дефект уже виявлено, яка ймовірність, що продукт виготовлено на лінії №2?

$$P(H_2 | D) = \frac{P(H_2)P(D | H_2)}{P(D)} = \frac{0.40 \cdot 0.05}{0.032} = \frac{0.020}{0.032} = 0.625.$$

Предметний висновок. Хоча лінія №2 виробляє лише 40% загального обсягу, через більший рівень дефектності вона є найбільш імовірним джерелом знайденого браку — 62.5%.

Завдання 4. Серійний контроль за схемою Бернуллі

За стабільного процесу ймовірність того, що випадково вибрана упаковка не відповідає нормативним вимогам, становить $p = 0.04$. Контролюють 12 незалежно відібраних упаковок.

$$n = 12, \quad p = 0.04, \quad q = 0.96.$$

1. Рівно одна дефектна упаковка.

$$P_{12}(1) = C_{12}^1(0.04)^1(0.96)^{11}.$$

$$P_{12}(1) = 12 \cdot 0.04 \cdot 0.96^{11} \approx 0.3064.$$

2. Жодної дефектної упаковки.

$$P_{12}(0) = 0.96^{12} \approx 0.6127.$$

3. Хоча б одна дефектна упаковка.

$$P(X \geq 1) = 1 - P(X = 0) = 1 - 0.6127 = 0.3873.$$

Предметний висновок. За рівня дефектності 4% контроль 12 одиниць дає приблизно 38.73% імовірності виявити хоча б одну дефектну одиницю. Відсутність дефекту в невеликій контрольній вибірці не означає повної відсутності дефектної продукції в потоці.

5. Виконання та перевірка в Microsoft Excel

Комп'ютерний етап виконується після ручного розрахунку: ручний розрахунок → Excel → зіставлення → висновок.

Розрахунок	Формула Excel
Комбінації	=COMBIN(10,3)
Розміщення	=PERMUT(10,3)
Рівно 1 дефект без повернення	=COMBIN(4,1)*COMBIN(16,2)/COMBIN(20,3)
Хоча б 1 дефект без повернення	=1-COMBIN(16,3)/COMBIN(20,3)
Бернуллі: рівно 1	=BINOM.DIST(1,12,0.04,FALSE)
Бернуллі: жодного	=BINOM.DIST(0,12,0.04,FALSE)
Бернуллі: хоча б 1	=1-BINOM.DIST(0,12,0.04,FALSE)

Для повної ймовірності та Байєса доцільно створити таблицю з колонками: лінія, $P(H_i)$, $P(D | H_i)$, добуток $P(H_i)P(D | H_i)$. Сума останнього стовпця дає $P(D)$, а частка відповідного внеску — байєсівську ймовірність.

6. Вимоги до оформлення звіту

Звіт повинен містити номер і назву роботи, мету, постановку задач, формули, обґрунтування вибору моделі, ручні розрахунки, результати Excel, таблицю зіставлення та короткі висновки. Недостатньо записати лише числову відповідь.

7. Контрольні питання

1. У чому полягає відмінність між перестановками, розміщеннями та комбінаціями?
2. Чому в задачах формування звичайної контрольної вибірки найчастіше використовують комбінації?
3. За яких умов можна використовувати класичне означення ймовірності?
4. Для чого використовують протилежну подію і чому вона зручна для задачі «хоча б один дефект»?
5. Коли для двох подій можна записати $P(A \cap B) = P(A)P(B)$?
6. Чому послідовний відбір із невеликої партії без повернення не можна автоматично розглядати як серію незалежних випробувань?
7. У чому різниця між формулою повної ймовірності та формулою Байєса?
8. Чому технологічна лінія з меншою часткою виробництва може бути найбільш імовірним джерелом знайденого дефекту?
9. Які умови повинні виконуватися для схеми Бернуллі?
10. Що зміниться в ймовірності виявлення хоча б одного дефекту при збільшенні обсягу вибірки?
11. Чи означає відсутність дефектних одиниць у невеликій вибірці, що їх немає в усій партії?
12. У чому принципова різниця між вибором без повернення з кінцевої партії та n незалежними випробуваннями Бернуллі?

Практична робота №2

Робота з вибірками, побудова гістограм та контрольних карт Шухарта

1. Мета роботи

Набути практичних навичок первинної статистичної обробки вибірових технологічних даних, їх графічного представлення та оцінювання стабільності технологічного процесу за допомогою контрольних карт Шухарта.

- формувати та аналізувати вибірку експериментальних або виробничих даних, обчислювати її основні числові характеристики;
- будувати інтервальний варіаційний ряд, гістограму та boxplot і характеризувати форму розподілу;
- виявляти потенційні викиди та коректно трактувати їх;
- будувати I–MR контрольні карти Шухарта для послідовних одиничних вимірювань;
- розрізнати контрольні межі статистичного процесу та нормативні межі технологічного допуску.

2. Короткі теоретичні відомості

2.1. Генеральна сукупність і вибірка

Генеральною сукупністю називають множину всіх об'єктів або результатів спостережень, щодо яких необхідно зробити статистичний висновок. Наприклад, усі упаковки певного продукту, виготовлені протягом зміни.

Повне дослідження генеральної сукупності часто неможливе або недоцільне, тому аналізують її частину — **вибірку**: $x_1, x_2, \dots, x_n$. Число n називають обсягом вибірки. Вибірka має бути репрезентативною.

2.2. Варіаційний ряд

Якщо значення вибірки впорядкувати за зростанням:

$$x_{(1)} \leq x_{(2)} \leq \dots \leq x_{(n)},$$

отримаємо варіаційний ряд. Він дає змогу визначити мінімум, максимум, розмах, побачити концентрацію спостережень та крайні значення.

2.3. Основні вибіркові характеристики

Вибіркове середнє:

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i.$$

Медіана – це величина, що ділить впорядковану вибірку на дві рівні за кількістю частини. Вона є стійкішою до крайніх значень, ніж середнє.

Розмах:

$$R = x_{\max} - x_{\min}.$$

Розмах є найпростішою характеристикою розсіювання, але залежить лише від двох крайніх спостережень.

Виправлена вибіркова дисперсія:

$$s^2 = \frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n - 1}.$$

Вибіркове стандартне відхилення:

$$s = \sqrt{s^2}.$$

Стандартне відхилення характеризує типовий масштаб розсіювання результатів навколо середнього.

Для додатних величин відносно розсіювання можна оцінювати коефіцієнтом варіації:

$$V = \frac{s}{\bar{x}} \cdot 100\%.$$

Він дає можливість порівнювати відносну мінливість показників, виражених у різних масштабах.

2.4. Квартилі та міжквартильний розмах

Перший квартиль Q_1 дорівнює значенню, що відповідає 25% вибірки, другий Q_2 є медіаною, третій Q_3 — 75% рівнем.

$$IQR = Q_3 - Q_1.$$

IQR характеризує розсіювання центральних 50% спостережень і є стійкішим до крайніх значень, ніж звичайний розмах.

2.5. Виявлення потенційних викидів

Для первинного аналізу використовується правило 1.5IQR:

$$L_f = Q_1 - 1.5IQR, \quad U_f = Q_3 + 1.5IQR.$$

Спостереження поза цими межами є **потенційними** викидами. Їх не можна автоматично вилучати. Необхідно перевірити можливі причини: помилка введення, несправність вимірювання, реальне порушення технологічного режиму або природна варіативність.

2.6. Інтервальний варіаційний ряд

Для великої кількості значень спостереження об'єднують в інтервали. Орієнтовну кількість інтервалів можна визначити за формулою Стерджеса:

$$k = 1 + 3.322 \lg n.$$

Ширина інтервалу:

$$h \approx \frac{R}{k}.$$

Для кожного інтервалу визначають абсолютну частоту n_j та відносну:

$$w_j = \frac{n_j}{n}.$$

Перевірка:

$$\sum_j n_j = n, \quad \sum_j w_j = 1.$$

2.7. Гістограма

Гістограма — графічне представлення розподілу кількісної величини за інтервалами. Вона дає змогу попередньо оцінити центр, розсіювання, симетрію, мультимодальність та можливі викиди.

Гістограма дає змогу попередньо оцінити:

- положення центра розподілу;
- розсіювання;
- симетричність або асиметричність;
- наявність декількох максимумів;
- можливі викиди;
- приблизну відповідність нормальному розподілу.

Висновок про нормальність не можна робити лише за виглядом гістограми. Формальна перевірка нормальності виконується в наступних практичних роботах.

2.8. Boxplot

Boxplot відображає Q_1 , медіану, Q_3 , IQR, крайні значення та потенційні викиди. Він особливо корисний для швидкого виявлення асиметрії та віддалених спостережень.

2.9. Статистичне управління технологічними процесами

Statistical Process Control (SPC) — сукупність статистичних методів для контролю стабільності процесу в часі.

Причини варіативності умовно поділяють на **звичайні** — дрібні постійні коливання, властиві процесу, та **спеціальні** — нетипові фактори, наприклад несправність дозатора, зміна партії сировини, помилка оператора або порушення температурного режиму.

Якщо діють лише звичайні причини, процес називають статистично керованим або стабільним.

2.10. Контрольна карта Шухарта

Контрольна карта відображає значення технологічного показника у порядку їх отримання та статистично визначені межі: центральну лінію CL , верхню UCL і нижню LCL .

2.11. Карта індивідуальних значень I–MR

Якщо в кожний момент часу отримують одне спостереження, використовують пару карт: I-chart — карта індивідуальних значень, MR-chart — карта рухомих розмахів.

Рухомий розмах:

$$MR_i = |x_i - x_{i-1}|, \quad i = 2, \dots, n.$$

Середній рухомий розмах:

$$\overline{MR} = \frac{1}{n-1} \sum_{i=2}^n MR_i.$$

Для рухомого розмаху двох сусідніх спостережень $d_2 = 1.128$, тому оцінка стандартного відхилення процесу:

$$\hat{\sigma} = \frac{\overline{MR}}{1.128}.$$

2.12. Контрольні межі I-карти

$$CL = \bar{x}.$$

$$UCL = \bar{x} + 3\hat{\sigma} = \bar{x} + 2.66\overline{MR}.$$

$$LCL = \bar{x} - 3\hat{\sigma} = \bar{x} - 2.66\overline{MR}.$$

2.13. Контрольні межі MR-карти

Для рухомого розмаху двох послідовних спостережень $D_3 = 0$, $D_4 = 3.267$:

$$CL_{MR} = \overline{MR}, \quad UCL_{MR} = 3.267\overline{MR}, \quad LCL_{MR} = 0.$$

2.14. Ознаки статистичної некерованості

Найпростіший сигнал спеціальної причини — вихід точки за UCL або LCL . Також підозрілими можуть бути довгі серії точок по один бік центральної лінії, монотонні тренди, періодичність або незвичайне групування. У цій роботі основною формальною ознакою вважається вихід за контрольні межі.

2.15. Контрольні межі і межі допуску — не одне й те саме

LCL, UCL визначаються за статистичними властивостями самого процесу та відповідають на питання про його стабільність. LSL, USL задаються нормативним документом або технологічною специфікацією та відповідають на питання про відповідність продукції вимогам.

2.16. Загальний алгоритм аналізу

дані $\rightarrow$ варіаційний ряд $\rightarrow$ описові характеристики $\rightarrow$ гістограма/boxplotчасова послідовність $\rightarrow$ I-MR $\rightarrow$ висновок

3. Порядок виконання роботи

1. Записати вихідну послідовність технологічних вимірювань у порядку їх отримання.
2. Визначити обсяг вибірки n .
3. Побудувати варіаційний ряд.
4. Обчислити $x_{\min}$, $x_{\max}$, R , середнє, медіану, s^2 , s , V , Q_1 , Q_3 , IQR .
5. За правилом $1.5IQR$ визначити потенційні викиди.
6. За Стерджемсом визначити орієнтовну кількість інтервалів і сформувати інтервальний ряд.
7. Побудувати гістограму та boxplot.
8. Повернутися до початкового часового порядку.
9. Для кожної пари сусідніх спостережень обчислити MR_i та середній MR .
10. Визначити CL , UCL , LCL для I-карти та побудувати її.
11. Визначити межі MR-карти та побудувати її.
12. Перевірити сигнали статистичної некерованості.
13. Повторити розрахунки у Microsoft Excel.
14. Сформулювати статистичний та предметний технологічний висновки.

4. Приклад виконання. Демонстраційний варіант 0

На фасувальній лінії харчового підприємства послідовно контролювали масу 20 упаковок продукту.

№ упаковки	Маса, г	№ упаковки	Маса, г
1	500.8	11	499.7
2	499.6	12	500.5
3	500.4	13	500.0
4	501.1	14	499.5
5	499.9	15	500.6
6	500.2	16	503.0
7	500.7	17	500.3
8	499.8	18	499.8
9	500.1	19	500.4
10	501.0	20	499.9

4.1. Обсяг вибірки та варіаційний ряд

$$n = 20.$$

Впорядковані значення: 499.5; 499.6; 499.7; 499.8; 499.8; 499.9; 499.9; 500.0; 500.1; 500.2; 500.3; 500.4; 500.4; 500.5; 500.6; 500.7; 500.8; 501.0; 501.1; 503.0.

$$x_{\min} = 499.5 \text{ г}, \quad x_{\max} = 503.0 \text{ г}.$$

$$R = 503.0 - 499.5 = 3.5 \text{ г}.$$

4.2. Вибіркове середнє

$$\bar{x} = \frac{1}{20} \sum_{i=1}^{20} x_i = 500.365 \text{ г}.$$

Після округлення: 500.37 г.

4.3. Медіана

Для парного $n=20$ медіана — середнє між 10-м і 11-м елементами варіаційного ряду:

$$Me = \frac{500.2 + 500.3}{2} = 500.25 \text{ г.}$$

Середнє дещо вище за медіану через велике значення 503.0 г.

4.4. Дисперсія, стандартне відхилення та V

$$s^2 \approx 0.5971 \text{ г}^2.$$

$$s = \sqrt{0.5971} \approx 0.7727 \text{ г.}$$

$$V = \frac{0.7727}{500.365} \cdot 100\% \approx 0.154\%.$$

Відносне розсіювання невелике, але це саме по собі не доводить статистичної стабільності процесу.

4.5. Квартилі та потенційні викиди

$$Q_1 = 499.875 \text{ г}, \quad Q_3 = 500.625 \text{ г.}$$

$$IQR = 500.625 - 499.875 = 0.750 \text{ г.}$$

$$L_f = 499.875 - 1.5 \cdot 0.750 = 498.750 \text{ г.}$$

$$U_f = 500.625 + 1.5 \cdot 0.750 = 501.750 \text{ г.}$$

Оскільки $503.0 > 501.750$, значення 503.0 г є потенційним викидом. Його не слід автоматично вилучати; потрібно перевірити можливі технологічні причини.

4.6. Інтервальний варіаційний ряд

$$k = 1 + 3.322 \lg 20 \approx 5.32.$$

Прийmemo 6 інтервалів. Орієнтовна ширина:

$$h = \frac{3.5}{6} \approx 0.58 \text{ г.}$$

Для зручності приймаємо $h = 0.6 \text{ г.}$

Інтервал маси, г	Частота n_j	Відносна частота w_j
$499.5 \leq x < 500.1$	8	0.40
$500.1 \leq x < 500.7$	7	0.35
$500.7 \leq x < 501.3$	4	0.20
$501.3 \leq x < 501.9$	0	0.00
$501.9 \leq x < 502.5$	0	0.00
$502.5 \leq x \leq 503.1$	1	0.05

Гістограма повинна показати концентрацію основної частини значень у діапазоні приблизно 499.5–501.3 г та окреме віддалене спостереження 503.0 г.

4.7. Розрахунок рухомих розмахів

Повертаємо дані до початкової часової послідовності. Наприклад:

$$MR_2 = |499.6 - 500.8| = 1.2.$$

$$MR_3 = |500.4 - 499.6| = 0.8.$$

Аналогічно обчислюються всі інші рухомі розмахи. Після підсумовування:

$$\overline{MR} \approx 0.9316 \text{ г.}$$

4.8. Оцінка стандартного відхилення процесу

$$\hat{\sigma} = \frac{0.9316}{1.128} \approx 0.8259 \text{ г.}$$

4.9. І-карта Шухарта

$$CL = \bar{x} = 500.365 \text{ г.}$$

$$UCL = 500.365 + 2.66 \cdot 0.9316 \approx 502.843 \text{ г.}$$

$$LCL = 500.365 - 2.66 \cdot 0.9316 \approx 497.887 \text{ г.}$$

Шістнадцяте спостереження:

$$x_{16} = 503.0 \text{ г.}$$

Оскільки:

$$503.0 > 502.843,$$

точка виходить за верхню контрольну межу.

4.10. MR-карта

$$CL_{MR} = 0.9316 \text{ г.}$$

$$UCL_{MR} = 3.267 \cdot 0.9316 \approx 3.043 \text{ г.}$$

$$LCL_{MR} = 0.$$

Найбільший рухомий розмах $MR_{17} = 2.7$ г не перевищує верхню контрольну межу.

4.11. Статистичний висновок

І-карта містить одне спостереження за верхньою контрольною межею. Це є статистичним сигналом можливої спеціальної причини, тому процес не можна вважати повністю статистично керованим протягом усього досліджуваного періоду.

4.12. Предметний технологічний висновок

Основна частина значень маси характеризується малим розсіюванням біля 500.4 г, але для упаковки №16 зафіксовано 503.0 г. Це значення одночасно є потенційним викидом за 1.5IQR та виходить за UCL І-карти. Потрібно перевірити роботу дозувального обладнання, умови фасування або вимірювання, а не просто вилучити результат.

5. Виконання та перевірка в Microsoft Excel

Показник	Функція / дія
n	=COUNT(B2:B21)
Середнє	=AVERAGE(B2:B21)
Медіана	=MEDIAN(B2:B21)
Мінімум	=MIN(B2:B21)
Максимум	=MAX(B2:B21)
Дисперсія	=VAR.S(B2:B21)
Стандартне відхилення	=STDEV.S(B2:B21)
Q1	=QUARTILE.INC(B2:B21,1)
Q3	=QUARTILE.INC(B2:B21,3)
MR	=ABS(B3-B2)

Гістограма: **Insert** → **Insert Statistic Chart** → **Histogram**. Boxplot: **Insert** → **Insert Statistic Chart** → **Box and Whisker**. Для І-карти створити стовпці «значення», CL, UCL, LCL і побудувати Line Chart. Для MR-карти — аналогічну таблицю рухомих розмахів.

6. Вимоги до оформлення звіту

Звіт повинен містити вихідні дані, варіаційний ряд, описові характеристики, квартилі та IQR, інтервальний ряд, гістограму, boxplot, таблицю MR, розрахунок меж I- та MR-карт, самі карти, результати Excel, статистичний і технологічний висновки.

7. Контрольні питання

1. У чому полягає відмінність між генеральною сукупністю та вибіркою?
2. Що означає репрезентативність і чому великий обсяг вибірки сам по собі не гарантує її?
3. Чим вибіркове середнє відрізняється від медіани?
4. Чому розмах є менш надійною характеристикою мінливості, ніж стандартне відхилення або IQR?
5. Що показує гістограма і чому за її виглядом не можна остаточно довести нормальність?
6. Що означає правило $1.5IQR$ і чи потрібно автоматично вилучати такі точки?
7. Чому перед побудовою гістограми дані можна впорядковувати, а перед контрольною картою — ні?
8. У чому принципова відмінність між гістограмою та контрольною картою?
9. Що означає вихід точки за контрольну межу?
10. Чому LCL/UCL не можна ототожнювати з LSL/USL?
11. Чи може процес бути статистично стабільним, але виробляти продукцію поза нормативами?
12. Чи гарантує відсутність точок за межами повну стабільність процесу?

Практична робота №3

Розрахунок описової статистики та довірчих інтервалів засобами Microsoft Excel і спеціалізованого статистичного ПЗ

1. Мета роботи

Набути практичних навичок статистичного опису результатів експериментальних досліджень та інтервального оцінювання параметрів генеральної сукупності.

- обчислювати та змістовно інтерпретувати основні характеристики положення і розсіювання вибірових даних;
- розрізняти стандартне відхилення та стандартну похибку середнього;
- будувати довірчий інтервал для генерального середнього за невідомої генеральної дисперсії;
- будувати довірчі інтервали для генеральної дисперсії та стандартного відхилення;
- виконувати ручний розрахунок, перевіряти результати у Microsoft Excel та спеціалізованому статистичному ПЗ;
- коректно інтерпретувати оцінки в контексті експериментальних досліджень у харчових технологіях.

2. Короткі теоретичні відомості

2.1. Статистичне оцінювання

Нехай досліджується кількісна характеристика технологічного процесу або харчового продукту: масова частка вологи, кислотність, концентрація компонента, маса продукту, тривалість технологічної операції тощо.

Генеральна сукупність характеризується параметрами μ — генеральним середнім, σ^2 — генеральною дисперсією, σ — генеральним стандартним відхиленням. Їх значення зазвичай невідомі і оцінюються за вибіркою $x_1, x_2, \dots, x_n$.

Розрізняють **точкові** та **інтервальні** статистичні оцінки.

2.2. Точкова оцінка

Точкова оцінка — одне числове значення, обчислене за вибіркою та використане для оцінювання невідомого параметра. Наприклад, $\bar{x}$ є оцінкою μ , а s^2 — оцінкою σ^2 .

2.3. Основні властивості статистичних оцінок

Незміщеність означає, що математичне сподівання оцінки дорівнює оцінюваному параметру.

Спроможність означає наближення оцінки до істинного параметра зі збільшенням обсягу вибірки.

Ефективність пов'язана з малою дисперсією оцінки серед інших незміщених оцінок.

2.4. Вибіркове середнє

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i.$$

Воно характеризує центральний рівень досліджуваної величини та є природною точковою оцінкою генерального середнього:

$$\hat{\mu} = \bar{x}.$$

2.5. Медіана та мода

Медіана — значення, яке ділить впорядковану вибірку на дві рівні за кількістю частини. Для парного обсягу:

$$Me = \frac{x_{(n/2)} + x_{(n/2+1)}}{2}.$$

Мода — значення, яке трапляється найчастіше.

2.6. Розмах, дисперсія та стандартне відхилення

$$R = x_{\max} - x_{\min}.$$

Виправлена вибіркова дисперсія:

$$s^2 = \frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n - 1}.$$

Число $\nu = n - 1$ називають числом ступенів свободи.

Стандартне відхилення:

$$s = \sqrt{s^2}.$$

Воно має ті самі одиниці вимірювання, що й вихідна величина, і характеризує типовий масштаб розсіювання окремих спостережень навколо середнього.

2.7. Коефіцієнт варіації

$$V = \frac{s}{\bar{x}} \cdot 100\%.$$

Це безрозмірна характеристика відносної мінливості. Її слід використовувати обережно, якщо середнє близьке до нуля або шкала не має змістовного абсолютного нуля.

2.8. Квартилі та міжквартильний розмах

Q_1 відповідає приблизно 25% вибірки, $Q_2 = Me$, Q_3 — 75% рівню.

$$IQR = Q_3 - Q_1.$$

2.9. Стандартна похибка середнього

Якщо багато разів формувати випадкові вибірки однакового обсягу, їх вибіркові середні будуть дещо відрізнятися. Мінливість вибіркового середнього характеризує стандартна похибка:

$$SE_{\bar{x}} = \frac{s}{\sqrt{n}}.$$

Принципова відмінність: s характеризує мінливість **окремих спостережень**, тоді як $SE_{\bar{x}}$ характеризує **точність оцінювання генерального середнього**.

$$\boxed{s \neq SE_{\bar{x}}}.$$

Оскільки $SE_{\bar{x}} = s/\sqrt{n}$, збільшення обсягу вибірки зменшує стандартну похибку. Щоб зменшити SE приблизно удвічі при незмінному s , n потрібно збільшити приблизно у 4 рази.

2.10. Інтервальна оцінка

Точкова оцінка не показує ступеня її невизначеності. Тому використовують довірчий інтервал:

$$\theta_L < \theta < \theta_U,$$

де θ — невідомий параметр, θ_L, θ_U — межі інтервалу.

2.11. Довірча ймовірність і рівень значущості

Довірчу ймовірність позначають

$$\gamma = 1 - \alpha.$$

Число α називають рівнем значущості.

Найчастіше використовують $\gamma = 0.95$, отже $\alpha = 0.05$. Такий інтервал називають 95-відсотковим довірчим інтервалом (Confidence Interval) та часто позначають 95% CI..

2.12. Правильна інтерпретація довірчого інтервалу

У прикладній мові часто говорять: «з довірою 95% генеральне середнє знаходиться в отриманому інтервалі». Строгіше: якби процедуру формування вибірки та побудови інтервалу багаторазово повторювали одним і тим самим методом, приблизно 95% побудованих інтервалів містили б істинне значення параметра.

Зокрема 95% СІ для середнього **не означає**, що 95% окремих спостережень знаходяться всередині нього.

2.13. Довірчий інтервал для генерального середнього

Якщо генеральна дисперсія невідома, використовується t-розподіл Стюдента:

$$\bar{x} - t_{1-\alpha/2; n-1} \frac{s}{\sqrt{n}} < \mu < \bar{x} + t_{1-\alpha/2; n-1} \frac{s}{\sqrt{n}}$$

Або коротко:

$$\mu \in \bar{x} \pm t_{1-\alpha/2; n-1} SE_{\bar{x}}.$$

Гранична похибка оцінювання середнього:

$$\Delta = t_{1-\alpha/2; n-1} \frac{s}{\sqrt{n}}.$$

2.14. Умови застосування t-інтервалу

- досліджувана величина є кількісною;
- вибірка є випадковою та репрезентативною;
- спостереження незалежні;
- генеральний розподіл приблизно нормальний, особливо для малих вибірок;
- відсутні необґрунтовані грубі помилки.

Для великих вибірок вимога нормальності стає менш жорсткою завдяки властивостям вибіркового середнього. У цій роботі формальна перевірка нормальності не є центральною; вона буде виконуватися у ПР №4.

2.15. Довірчий інтервал для генеральної дисперсії

Якщо генеральна сукупність має нормальний розподіл, статистика $(n-1)s^2/\sigma^2$ має розподіл χ^2 з $n-1$ ступенями свободи. Тоді:

$$\frac{(n-1)s^2}{\chi_{1-\alpha/2; n-1}^2} < \sigma^2 < \frac{(n-1)s^2}{\chi_{\alpha/2; n-1}^2}.$$

На відміну від інтервалу для середнього, цей інтервал зазвичай несиметричний.

2.16. Довірчий інтервал для стандартного відхилення

Межі отримують добуванням квадратного кореня з меж інтервалу для дисперсії:

$$\sqrt{\sigma_L^2} < \sigma < \sqrt{\sigma_U^2}.$$

Класичний chi-square інтервал для дисперсії суттєво залежить від припущення про нормальність генеральної сукупності.

2.17. Що впливає на ширину довірчого інтервалу

- Чим більше s , тим ширший інтервал.
- Чим більше n , тим менша стандартна похибка і тим вужчий інтервал.
- Чим вищий рівень довіри, тим ширший інтервал.

2.18. Загальний алгоритм статистичного оцінювання

вибірка → описова статистика → s → SE → рівень довіри → t -CI для μ → χ^2 -CI для σ^2 → інтерпретація

3. Порядок виконання роботи

1. Записати вихідну вибірку досліджуваного технологічного показника.
2. Визначити обсяг вибірки n та побудувати варіаційний ряд.
3. Обчислити середнє, медіану, моду, мінімум, максимум, R , Q_1 , Q_3 , IQR .
4. Обчислити виправлену вибіркиму дисперсію s^2 .
5. Обчислити вибіркове стандартне відхилення s .
6. Обчислити коефіцієнт варіації V .
7. Обчислити стандартну похибку середнього SE .
8. Задати $\gamma=0.95$, $\alpha=0.05$.
9. Визначити $\nu=n-1$ та критичне t .
10. Визначити граничну похибку Δ .
11. Побудувати 95% CI для μ .
12. Визначити необхідні квантилі χ^2 .
13. Побудувати 95% CI для σ^2 та σ .
14. Повторити розрахунки у Microsoft Excel.
15. Перевірити основні результати у спеціалізованому ПЗ.
16. Зіставити результати та сформулювати статистичний і предметний висновки.

4. Приклад виконання. Демонстраційний варіант 0

Під час лабораторного контролю партії пшеничного борошна визначено масову частку вологи у 12 незалежно відібраних зразках.

№ зразка	Масова частка вологи, %	№ зразка	Масова частка вологи, %
1	13.7	7	13.9
2	13.9	8	14.2
3	14.1	9	13.8
4	13.8	10	14.1
5	14.0	11	13.9
6	13.6	12	14.0

4.1. Обсяг вибірки та варіаційний ряд

$$n = 12.$$

Впорядковані дані: 13.6; 13.7; 13.8; 13.8; 13.9; 13.9; 13.9; 14.0; 14.0; 14.1; 14.1; 14.2.

$$x_{\min} = 13.6\%, \quad x_{\max} = 14.2\%, \quad R = 0.6.$$

4.2. Вибіркове середнє

Сума спостережень:

$$\sum_{i=1}^{12} x_i = 167.0.$$

Тому:

$$\bar{x} = \frac{167.0}{12} = 13.9167\%.$$

4.3. Медіана та мода

Для парного $n=12$ медіана визначається за 6-м і 7-м елементами варіаційного ряду. Обидва дорівнюють 13.9:

$$Me = 13.9\%.$$

Значення 13.9 трапляється тричі — частіше за інші, тому:

$$Mo = 13.9\%.$$

4.4. Квартилі

За алгоритмом, який відповідає функції QUARTILE.INC у Microsoft Excel:

$$Q_1 = 13.800\%, \quad Q_3 = 14.025\%.$$

$$IQR = 14.025 - 13.800 = 0.225.$$

4.5. Вибіркова дисперсія

Розраховуємо суму квадратів відхилень від середнього:

$$\sum (x_i - \bar{x})^2 \approx 0.33667.$$

Число ступенів свободи:

$$n - 1 = 11.$$

Тому:

$$s^2 = \frac{0.33667}{11} \approx 0.030606.$$

4.6. Стандартне відхилення

$$s = \sqrt{0.030606} \approx 0.17495.$$

Типовий масштаб розсіювання окремих результатів відносно середнього становить приблизно 0.175 відсоткового пункту.

4.7. Коефіцієнт варіації

$$V = \frac{0.17495}{13.9167} \cdot 100\% \approx 1.257\%.$$

4.8. Стандартна похибка середнього

$$SE_{\bar{x}} = \frac{0.17495}{\sqrt{12}} \approx 0.05050.$$

Не слід стверджувати, що «похибка кожного вимірювання дорівнює 0.0505%». Це характеристика невизначеності оцінки середнього.

4.9. Критичне значення t

$$\gamma = 0.95, \quad \alpha = 0.05, \quad \nu = 11.$$

$$t_{0.975;11} \approx 2.201.$$

4.10. Гранична похибка

$$\Delta = t_{0.975;11} SE_{\bar{x}} = 2.201 \cdot 0.05050 \approx 0.11116.$$

4.11. Довірчий інтервал для генерального середнього

$$\mu_L = 13.9167 - 0.1112 = 13.8055.$$

$$\mu_U = 13.9167 + 0.1112 = 14.0278.$$

$$95\% CI_{\mu} = [13.8055; 14.0278]\%.$$

Практично: середню масову частку вологи у генеральній сукупності оцінюємо приблизно як 13.92%, а інтервальна оцінка при рівні довіри 95% становить приблизно 13.81–14.03%.

4.12. Довірчий інтервал для генеральної дисперсії

Для $v=11$ і $\alpha=0.05$:

$$\chi_{0.025;11}^2 \approx 3.8157, \quad \chi_{0.975;11}^2 \approx 21.9200.$$

Маємо:

$$(n-1)s^2 = 11 \cdot 0.030606 \approx 0.33667.$$

Нижня межа:

$$\sigma_L^2 = \frac{0.33667}{21.9200} \approx 0.01536.$$

Верхня межа:

$$\sigma_U^2 = \frac{0.33667}{3.8157} \approx 0.08823.$$

$$95\% CI_{\sigma^2} = [0.01536; 0.08823].$$

4.13. Довірчий інтервал для стандартного відхилення

$$\sigma_L = \sqrt{0.01536} \approx 0.1239.$$

$$\sigma_U = \sqrt{0.08823} \approx 0.2970.$$

$$95\% CI_{\sigma} = [0.1239; 0.2970].$$

Інтервал для стандартного відхилення несиметричний відносно вибіркового значення $s \approx 0.175$.

4.14. Підсумкова таблиця

Характеристика	Значення
n	12
Середнє, %	13.9167
Медіана, %	13.9000
Мода, %	13.9000
Мінімум, %	13.6000
Максимум, %	14.2000
Розмах	0.6000
Q1, %	13.8000
Q3, %	14.0250
IQR	0.2250
s^2	0.03061
s	0.17495
V, %	1.257
SE	0.05050
95% CI для μ , %	13.8055–14.0278
95% CI для σ^2	0.01536–0.08823
95% CI для σ	0.1239–0.2970

4.15. Статистичний висновок

За вибіркою з 12 спостережень середня масова частка вологи становить 13.917%, стандартне відхилення — 0.175 відсоткового пункту, стандартна похибка середнього — 0.0505. 95% довірчий інтервал для генерального середнього становить 13.806–14.028%.

4.16. Предметний технологічний висновок

Вибірka характеризується невеликим відносним розсіюванням: $V \approx 1.26\%$. Середню масову частку вологи партії можна оцінити приблизно як 13.92%, але разом із точковою оцінкою доцільно повідомляти 95% CI 13.81–14.03%, оскільки він відображає невизначеність вибіркового оцінювання.

5. Виконання та перевірка в Microsoft Excel

Показник	Формула Excel
n	=COUNT(B2:B13)
Середнє	=AVERAGE(B2:B13)
Медіана	=MEDIAN(B2:B13)
Мода	=MODE.SNGL(B2:B13)
Q1	=QUARTILE.INC(B2:B13,1)
Q3	=QUARTILE.INC(B2:B13,3)
Дисперсія	=VAR.S(B2:B13)
Стандартне відхилення	=STDEV.S(B2:B13)
SE	=STDEV.S(B2:B13)/SQRT(COUNT(B2:B13))
t критичне	=T.INV.2T(0.05,COUNT(B2:B13)-1)
Гранична похибка	=CONFIDENCE.T(0.05,STDEV.S(B2:B13),COUNT(B2:B13))

Для нижньої межі СІ дисперсії можна використати $=(n-1)*VAR.S(range)/CHISQ.INV(0.975,n-1)$, для верхньої — ту саму формулу з імовірністю 0.025. Для σ беруть квадратний корінь.

Додатково: **Data** → **Data Analysis** → **Descriptive Statistics**. Спеціалізоване ПЗ використовується для незалежної перевірки та демонстрації професійного статистичного середовища.

6. Вимоги до оформлення звіту

Звіт має містити вихідні дані, варіаційний ряд, середнє, медіану, моду, мінімум, максимум, квартилі, IQR, ручні s^2 , s , V , SE , рівень довіри, t -критичне, 95% СІ для μ , σ^2 , σ , результати Excel/ПЗ, таблицю зіставлення, статистичний і предметний висновки.

7. Контрольні питання

1. У чому різниця між точковою та інтервальною оцінкою?
2. Чому у формулі виправленої вибіркової дисперсії використовується $n-1$, а не n ?
3. У чому різниця між стандартним відхиленням s і стандартною похибкою SE ?
4. Що станеться зі стандартною похибкою, якщо n збільшити у 4 рази при незмінному s ?
5. Чому при невідомій дисперсії для СІ середнього використовується t -розподіл?
6. Що означає 95% довірчий інтервал і чому неправильно трактувати його як інтервал для 95% окремих спостережень?
7. Що станеться з СІ при збільшенні обсягу вибірки?
8. Що станеться з шириною СІ при переході від 95% до 99% довіри?
9. Чому СІ для генеральної дисперсії несиметричний?
10. Чому χ^2 -square інтервал для дисперсії не слід механічно застосовувати до сильно ненормальних даних?
11. Чому $VAR.S/STDEV.S$ використовуються для вибірки, а $VAR.P/STDEV.P$ — для іншої ситуації?
12. Чи гарантує малий коефіцієнт варіації високу точність оцінювання генерального середнього?

Практична робота №4

Комплексний порівняльний аналіз двох технологічних процесів: перевірка нормальності, F-критерій рівності дисперсій та t-критерій Стюдента

1. Мета роботи

Набути практичних навичок комплексного статистичного порівняння двох незалежних технологічних процесів або двох варіантів технології.

- формулювати нульову та альтернативну статистичні гіпотези для двох незалежних вибірок;
- перевіряти припущення про нормальність розподілу досліджуваного показника;
- перевіряти гіпотезу про рівність генеральних дисперсій за F-критерієм Фішера;
- обґрунтовано вибирати між класичним t-критерієм Стюдента та t-критерієм Велча;
- перевіряти гіпотезу про рівність генеральних середніх;
- будувати довірчий інтервал для різниці середніх та оцінювати Cohen's d;
- розрізняти статистичну та технологічну значущість.

2. Короткі теоретичні відомості

2.1. Постановка задачі порівняння двох технологічних процесів

У технологічних дослідженнях часто необхідно встановити, чи відрізняються результати двох технологій: двох режимів сушіння, рецептур, способів термічного оброблення, типів пакування, режимів ферментації тощо.

Нехай для першого процесу отримано незалежну вибірку $x_1, x_2, \dots, x_{n_1}$, для другого — $y_1, y_2, \dots, y_{n_2}$. Генеральні середні: μ_1, μ_2 , дисперсії: σ_1^2, σ_2^2 .

Основна задача — встановити, чи є спостережувана різниця між вибірковими середніми наслідком реальної відмінності технологій, чи її можна пояснити випадковою варіативністю.

2.2. Незалежні вибірки

Дві вибірки є незалежними, якщо результати однієї групи не пов'язані попарно з результатами іншої. Якщо ті самі об'єкти вимірюють до та після впливу, це парні дані і потрібен інший варіант t-тесту.

2.3. Загальна схема перевірки статистичної гіпотези

1. Сформулювати H_0 .
2. Сформулювати H_1 .
3. Задати рівень значущості α .
4. Вибрати статистичний критерій.
5. Обчислити статистику критерію.
6. Визначити критичне значення або p-value.
7. Прийняти статистичне рішення.
8. Інтерпретувати результат у контексті дослідження.

У цій роботі приймаємо:

$$\alpha = 0.05.$$

2.4. Нульова та альтернативна гіпотези

$$H_0: \mu_1 = \mu_2.$$

$$H_1: \mu_1 \neq \mu_2.$$

Використовується двостороння перевірка. Односторонню альтернативу доцільно задавати лише тоді, коли напрям ефекту був обґрунтований до аналізу даних.

2.5. p-value

p-value — ймовірність отримати результат, не менш суперечливий із H_0 , ніж фактично спостережуваний, за умови правильності H_0 . Якщо $p \leq \alpha$, H_0 відхиляють; якщо $p > \alpha$, підстав для відхилення H_0 недостатньо.

Результат $p > 0.05$ не доводить правильності H_0 . Результат $p = 0.004$ не означає, що ймовірність правильності H_0 дорівнює 0.4%.

2.6. Передумови параметричного порівняння

- кількісний характер досліджуваної величини;
- незалежність спостережень;
- відсутність необґрунтованих грубих помилок;
- приблизна нормальність розподілу в кожній групі;
- для класичного pooled t-test — приблизна рівність дисперсій.

2.7. Первинний графічний аналіз

Перед формальними тестами необхідно дослідити дані графічно. Доцільно побудувати гістограми, boxplot та Q–Q plots. Вони дозволяють виявити сильну асиметрію, потенційні викиди, мультимодальність та великі відмінності в розсіюванні.

2.8. Перевірка нормальності за Шапіро–Вілком

Для кожної вибірки окремо перевіряють:

H_0 : дані узгоджуються з нормальним розподілом.

H_1 : дані не узгоджуються з нормальним розподілом.

Статистика критерію позначається W . Її коефіцієнти вручну в цій роботі не розраховуються — Shapiro–Wilk виконується в Real Statistics або Jamovi.

Якщо $p > 0.05$, немає достатніх підстав відхилити припущення про нормальність. Це не означає, що нормальність доведено. Для малих вибірок критерій може мати недостатню потужність, тому результат потрібно розглядати разом із графіками.

2.9. Q–Q plot

Q–Q plot порівнює емпіричні значення вибірки з теоретичними квантилями нормального розподілу. За приближної нормальності точки розташовуються близько до прямої. Систематична кривизна або сильні відхилення крайніх точок можуть свідчити про асиметрію, важкі хвости або викиди.

2.10. Перевірка рівності дисперсій

$$H_0: \sigma_1^2 = \sigma_2^2, \quad H_1: \sigma_1^2 \neq \sigma_2^2.$$

При нормальності генеральних сукупностей використовується F-критерій Фішера. Для зручності більшу дисперсію ставлять у чисельник:

$$F = \frac{s_{\max}^2}{s_{\min}^2} \geq 1.$$

Ступені свободи відповідають дисперсіям у чисельнику та знаменнику. Для двосторонньої перевірки при $\alpha=0.05$ порівнюють F_{emp} з верхнім критичним значенням $F_{1-\alpha/2; v_1, v_2}$.

2.11. Обмеження F-критерію

F-критерій є чутливим до відхилень від нормального розподілу. Тому його не слід механічно використовувати для явно ненормальних даних. У сучасному прикладному аналізі часто застосовують стійкіший критерій Левена. У цій роботі F-критерій зберігається як класичний метод, узгоджений із теоретичною частиною курсу.

2.12. Вибір t-критерію

Якщо дисперсії можна вважати рівними, використовують класичний t-критерій Стюдента для двох незалежних вибірок із pooled variance. Якщо дисперсії істотно різняться, використовують t-критерій Велча.

2.13. Об'єднана оцінка дисперсії

$$s_p^2 = \frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2}.$$
$$s_p = \sqrt{s_p^2}.$$

2.14. t-критерій Стюдента для незалежних вибірок

$$SE_{\bar{x}_1 - \bar{x}_2} = s_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}.$$
$$t = \frac{\bar{x}_1 - \bar{x}_2}{s_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}}.$$
$$\nu = n_1 + n_2 - 2.$$

Для двосторонньої перевірки H_0 відхиляють, якщо $|t_{emp}| > t_{crit}$ або $p < \alpha$.

2.15. t-критерій Велча

$$t = \frac{\bar{x}_1 - \bar{x}_2}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}}.$$

Число ступенів свободи оцінюється за формулою Велча–Саттертвейта:

$$\nu \approx \frac{\left(\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}\right)^2}{\frac{(s_1^2/n_1)^2}{n_1 - 1} + \frac{(s_2^2/n_2)^2}{n_2 - 1}}.$$

2.16. Довірчий інтервал для різниці середніх

Недостатньо лише встановити статистичну значущість. Необхідно оцінити величину різниці:

$$(\mu_1 - \mu_2) \in (\bar{x}_1 - \bar{x}_2) \pm t_{1-\alpha/2; \nu} SE_{\bar{x}_1 - \bar{x}_2}.$$

Якщо двосторонній 95% CI не містить нуль, це узгоджується зі статистично значущою різницею на рівні $\alpha=0.05$.

2.17. Розмір ефекту Cohen's d

p-value відповідає на питання про статистичні підстави для відмінності, але не характеризує безпосередньо її величину. За приблизно рівних дисперсій:

$$d = \frac{\bar{x}_1 - \bar{x}_2}{s_p}.$$

Орієнтовно $|d| \approx 0.2$ — малий ефект, 0.5 — середній, 0.8 і більше — великий. У технологічному дослідженні вирішальне значення має предметний зміст різниці.

2.18. Статистична і технологічна значущість

Можлива ситуація $p < 0.05$, але різниця технологічно мізерна. І навпаки, технологічно важлива різниця у малій вибірці може не досягти статистичної значущості через недостатню потужність.

Тому остаточний висновок повинен містити p -value, величину різниці, CI, effect size і технологічну інтерпретацію.

2.19. Алгоритм вибору методу

графічний аналіз → Shapiro–Wilk → перевірка дисперсій → Student/Welch → p → CI → d → висновок

3. Порядок виконання роботи

1. Записати результати двох незалежних вибірок.
2. Сформулювати технологічну задачу порівняння.
3. Для кожної вибірки обчислити n , середнє, медіану, дисперсію, стандартне відхилення, мінімум і максимум.
4. Побудувати boxplot, за можливості гістограму та Q–Q plot.
5. Для кожної групи сформулювати гіпотези нормальності і виконати Shapiro–Wilk.
6. Якщо параметричний підхід обґрунтований, сформулювати гіпотези про рівність дисперсій.
7. Вручну обчислити F-статистику та визначити критичне значення або p -value.
8. Прийняти рішення щодо рівності дисперсій.
9. Вибрати Student або Welch t -test.
10. Сформулювати H_0 та H_1 щодо середніх.
11. Вручну обчислити t -статистику і ступені свободи.
12. Визначити t_{crit} та p -value.
13. Побудувати 95% CI різниці середніх.
14. Обчислити Cohen's d .
15. Перевірити результати в Microsoft Excel та Real Statistics/Jamovi.
16. Сформулювати статистичний та предметний технологічний висновки.

4. Приклад виконання. Демонстраційний варіант 0

Досліджується вплив двох режимів сушіння фруктового продукту на кінцеву масову частку вологи. За кожним режимом отримано по 10 незалежних спостережень.

№	Режим А, %	Режим В, %
1	12.4	12.2
2	12.7	12.4
3	12.5	12.3
4	12.8	12.5
5	12.6	12.1
6	12.9	12.6
7	12.3	12.3
8	12.7	12.4
9	12.5	12.2
10	12.6	12.5

Приймаємо $\alpha = 0.05$.

4.1. Описова статистика

$$n_A = n_B = 10.$$

Для режиму А сума дорівнює 126.0:

$$\bar{x}_A = \frac{126.0}{10} = 12.60\%.$$

Для режиму В сума дорівнює 123.5:

$$\bar{x}_B = \frac{123.5}{10} = 12.35\%.$$

Різниця вибірових середніх:

$$\bar{x}_A - \bar{x}_B = 0.25 \text{ відсоткового пункту.}$$

4.2. Дисперсія першої вибірки

Для режиму А:

$$\sum (x_i - \bar{x}_A)^2 = 0.3000.$$

$$s_A^2 = \frac{0.3000}{9} = 0.03333.$$

$$s_A = \sqrt{0.03333} = 0.1826.$$

4.3. Дисперсія другої вибірки

Для режиму В:

$$\sum (y_i - \bar{x}_B)^2 = 0.2250.$$

$$s_B^2 = \frac{0.2250}{9} = 0.02500.$$

$$s_B = \sqrt{0.02500} = 0.1581.$$

4.4. Перевірка нормальності

Для кожного режиму перевіряємо H_0 : вибірка узгоджується з нормальним розподілом. За Shapiro–Wilk програмно отримано:

$$W_A \approx 0.9837, \quad p_A \approx 0.9819.$$

$$W_B \approx 0.9657, \quad p_B \approx 0.8486.$$

Обидва p-value більші 0.05, тому немає достатніх підстав відхилити припущення про нормальність. Це не є доказом абсолютної нормальності, але разом із графічним аналізом підтримує параметричний підхід.

4.5. Перевірка рівності дисперсій за F-критерієм

$$H_0: \sigma_A^2 = \sigma_B^2, \quad H_1: \sigma_A^2 \neq \sigma_B^2.$$

Більша дисперсія знаходиться у групі А:

$$F_{emp} = \frac{0.03333}{0.02500} = 1.3333.$$

$$v_1 = 9, \quad v_2 = 9.$$

Для двосторонньої перевірки при $\alpha=0.05$ верхнє критичне значення:

$$F_{0.975;9,9} \approx 4.026.$$

Порівнюємо:

$$1.333 < 4.026.$$

Отже, немає підстав відхиляти гіпотезу про рівність дисперсій. Excel дає приблизно:

$$p_F \approx 0.6752 > 0.05.$$

Для подальшого аналізу використовуємо класичний t-критерій Стюдента з pooled variance.

4.6. Об'єднана дисперсія

$$s_p^2 = \frac{9 \cdot 0.03333 + 9 \cdot 0.02500}{18}.$$

$$s_p^2 = \frac{0.3000 + 0.2250}{18} = 0.029167.$$

$$s_p = \sqrt{0.029167} \approx 0.17078.$$

4.7. Стандартна похибка різниці середніх

$$SE = s_p \sqrt{\frac{1}{10} + \frac{1}{10}}.$$

$$SE = 0.17078\sqrt{0.2} \approx 0.07638.$$

4.8. Розрахунок t-статистики

$$H_0: \mu_A = \mu_B, \quad H_1: \mu_A \neq \mu_B.$$

$$t_{emp} = \frac{12.60 - 12.35}{0.07638} \approx 3.273.$$

Ступені свободи:

$$\nu = 10 + 10 - 2 = 18.$$

4.9. Критичне значення і p-value

$$t_{0.975;18} \approx 2.101.$$

Оскільки:

$$|3.273| > 2.101,$$

Но відхиляємо. Програмний результат:

$$p \approx 0.00422 < 0.05.$$

Отже, середні значення для двох режимів статистично значущо відрізняються.

4.10. 95% довірчий інтервал для різниці середніх

Оцінена різниця: 0.25. Гранична похибка:

$$\Delta = 2.1009 \cdot 0.07638 \approx 0.16046.$$

Нижня межа:

$$0.25 - 0.16046 = 0.08954.$$

Верхня межа:

$$0.25 + 0.16046 = 0.41046.$$

$$95\% CI_{\mu_A - \mu_B} = [0.0895; 0.4105].$$

Інтервал не містить 0, що узгоджується з результатом t-тесту.

4.11. Cohen's d

$$d = \frac{0.25}{0.17078} \approx 1.464.$$

За орієнтовною шкалою це відповідає великому стандартизованому ефекту.

4.12. Статистичний висновок

Перевірка Shapiro–Wilk не виявила статистично значущих відхилень від нормальності. F-критерій не дав підстав вважати дисперсії різними. Тому застосовано класичний t-критерій Стьюдента. Отримано $t(18) = 3.273$, $p \approx 0.0042$. Різниця середніх 0.25 відсоткового пункту, 95% CI [0.090; 0.410], Cohen's $d \approx 1.46$.

4.13. Предметний технологічний висновок

Режим В забезпечує в середньому нижчу кінцеву масову частку вологи: 12.35% проти 12.60%. Спостережувана різниця статистично значуща. Водночас остаточне рішення про технологічну перевагу режиму В має враховувати нормативний діапазон вологості, енергетичні витрати, тривалість сушіння, органолептичні властивості та стабільність продукту під час зберігання.

5. Виконання та перевірка в Microsoft Excel

Описова статистика: AVERAGE, VAR.S, STDEV.S. Shapiro–Wilk виконується у Real Statistics або Jamovi.

Процедура	Excel
F-test p-value	=F.TEST(A2:A11,B2:B11)
Критичне F	=F.INV.RT(0.025,9,9)
Student, рівні дисперсії	=T.TEST(A2:A11,B2:B11,2,2)
Welch, нерівні дисперсії	=T.TEST(A2:A11,B2:B11,2,3)
t критичне	=T.INV.2T(0.05,18)

У Data Analysis ToolPak також доступні F-Test Two-Sample for Variances, t-Test: Two-Sample Assuming Equal Variances і t-Test: Two-Sample Assuming Unequal Variances. 95% CI різниці та Cohen's d обчислюють окремо.

6. Вимоги до оформлення звіту

Звіт повинен містити постановку задачі, вихідні дані, описову статистику, графічний аналіз, Shapiro–Wilk, ручний F-test, критичне F або p-value, обґрунтування вибору Student/Welch, ручний t-test, df, tcrit, p-value, 95% CI різниці, Cohen's d, результати Excel/ПЗ, статистичний і технологічний висновки.

7. Контрольні питання

1. Які основні умови повинні виконуватися для застосування t-критерію Стюдента до двох незалежних вибірок?
2. Що перевіряє критерій Шапіро–Вілکا і як правильно трактувати $p > 0.05$?
3. Чому не можна стверджувати, що $p > 0.05$ доводить нормальність?
4. Навіщо поряд із тестом нормальності аналізувати гістограму, boxplot або Q–Q plot?
5. Яку нульову гіпотезу перевіряє F-критерій?
6. Чому F-критерій небажано використовувати для сильно ненормальних даних?
7. У чому різниця між класичним Student t-test і Welch t-test?
8. Що слід зробити, якщо дисперсії істотно нерівні?
9. Чому при двосторонній альтернативі порівнюють $|t_{emp}|$ із критичним значенням?
10. Що означає $p = 0.004$ і чи є це ймовірністю правильності H_0 ?
11. Що показує 95% довірчий інтервал для різниці середніх?
12. Навіщо поряд із p-value обчислювати Cohen's d і чому статистична значущість не дорівнює технологічній важливості?

Практична робота №5

Органолептичний аналіз: обробка дегустаційних оцінок непараметричними критеріями Манна–Уїтні та Уїлкоксона

1. Мета роботи

Набути практичних навичок статистичного аналізу органолептичних та інших рангових даних за допомогою непараметричних критеріїв.

- визначати ситуації, у яких застосування параметричних критеріїв є небажаним або некоректним;
- розрізняти незалежні та парні вибірки й обґрунтовано вибирати між критеріями Манна–Уїтні та Уїлкоксона;
- виконувати ранжування спостережень та вручну обчислювати статистики U і T;
- інтерпретувати p-value і розмір непараметричного ефекту;
- формулювати статистичний і предметний висновок за результатами органолептичного дослідження;
- перевіряти ручні розрахунки в Microsoft Excel, Real Statistics або Jamovi.

2. Короткі теоретичні відомості

2.1. Чому виникає потреба в непараметричних критеріях

Параметричні критерії, зокрема t-критерій Стюдента, ґрунтуються на певних припущеннях щодо статистичної моделі. У реальних технологічних дослідженнях ці припущення можуть не виконуватися.

Особливо це характерно для малих вибірок, асиметричних розподілів, даних із викидами, порядкових шкал та балових органолептичних оцінок. У таких ситуаціях використовують **непараметричні статистичні критерії**, які часто працюють не безпосередньо з числовими значеннями, а з їх рангами.

2.2. Органолептичні дані

Органолептичний аналіз ґрунтується на оцінюванні властивостей продукту за допомогою органів чуття. Можуть оцінюватися смак, аромат, колір, консистенція, зовнішній вигляд, загальна прийнятність продукту.

Результат часто задається балами, наприклад 1–9. Такі оцінки мають природний порядок, але різниця між 8 і 7 балами не обов'язково має той самий зміст, що й різниця між 4 і 3 балами. Тому балові органолептичні дані часто доцільно аналізувати ранговими методами.

2.3. Ранг

Рангом називають порядковий номер спостереження у варіаційному ряду. Якщо декілька спостережень мають однакові значення, їм присвоюють середній ранг.

Наприклад, для значень 5, 6, 6, 8 два значення 6 займають позиції 2 і 3, тому кожному присвоюється середній ранг:

$$\frac{2 + 3}{2} = 2.5.$$

Отже, ранги: 1; 2.5; 2.5; 4.

2.4. Параметричний і непараметричний підходи

Для двох незалежних груп типовими альтернативами є:

t-критерій Стюдента ↔ Mann–Whitney U

Для двох залежних вибірок:

Вибір критерію залежить не лише від тесту нормальності, а й від типу шкали, дизайну експерименту, обсягу вибірки, характеру розподілу та наявності викидів.

2.5. Незалежні та залежні вибірки

Незалежні вибірки: спостереження однієї групи не мають попарного зв'язку зі спостереженнями іншої. Наприклад, продукт А оцінює одна група дегустаторів, продукт В — інша незалежна група.

Парні вибірки: кожне спостереження першої вибірки має природну пару у другій. Наприклад, той самий дегустатор оцінює продукт до і після зміни рецептури.

2.6. Критерій Манна–Уїтні

Критерій Манна–Уїтні призначений для порівняння двох незалежних вибірок. Нульова гіпотеза у загальному вигляді означає відсутність систематичного зсуву між розподілами двох груп.

Якщо розподіли мають приблизно однакову форму і відрізняються переважно зсувом, результат можна інтерпретувати як порівняння їх центрального положення. Проте Mann–Whitney не є буквально «тестом рівності медіан» у будь-якій ситуації.

2.7. Алгоритм критерію Манна–Уїтні

1. Об'єднати дві вибірки.
2. Упорядкувати всі значення за зростанням.
3. Присвоїти ранги.
4. Для однакових значень використати середні ранги.
5. Знайти суму рангів першої групи R_1 .
6. Знайти суму рангів другої групи R_2 .
7. Обчислити U_1 та U_2 .
8. Для ручного порівняння використати меншу статистику U .

$$U_1 = R_1 - \frac{n_1(n_1 + 1)}{2}, \quad U_2 = R_2 - \frac{n_2(n_2 + 1)}{2}.$$

$$U_1 + U_2 = n_1 n_2.$$

$$U = \min(U_1, U_2).$$

2.8. Прийняття рішення за Манна–Уїтні

Для малих вибірок можна використовувати критичне значення U_{crit} . Якщо $U \leq U_{crit}$, H_0 відхиляють. У програмному аналізі частіше використовують p-value: при $p \leq 0.05$ результат статистично значущий.

2.9. Однакові ранги — ties

Для органолептичних оцінок однакові бали трапляються дуже часто. При ручному ранжуванні їм присвоюють середні ранги. Програмне забезпечення при розрахунку p-value повинно враховувати поправку на ties.

2.10. Критерій знакових рангів Уїлкоксона

Wilcoxon signed-rank застосовується до парних спостережень. Для кожної пари:

$$d_i = y_i - x_i.$$

Нульова гіпотеза: систематичного зсуву парних різниць немає. Для приблизно симетричного розподілу різниць це можна трактувати як $Me_d = 0$.

2.11. Умови застосування Уїлкоксона

- спостереження є парними;

- пари незалежні одна від одної;
- різниці можна змістовно ранжувати за абсолютною величиною;
- для класичної інтерпретації зсуву бажана приблизна симетрія розподілу різниць.

Якщо симетричність різниць грубо порушена, може бути доцільним простіший критерій знаків.

2.12. Алгоритм критерію Уїлкоксона

1. Обчислити d_i .
2. Вилучити різниці $d_i=0$.
3. Обчислити $|d_i|$.
4. Ранжувати абсолютні різниці.
5. Повернути кожному рангу знак початкової різниці.
6. Знайти суму позитивних рангів T_+ .
7. Знайти суму від'ємних рангів T_- .
8. Визначити $T=\min(T_+, T_-)$.

$$T_+ = \sum R_i^+, \quad T_- = \sum R_i^-, \quad \boxed{T = \min(T_+, T_-)}.$$

2.13. Нульові різниці

Якщо $d_i = 0$, така пара не бере участі в ранжуванні. Тому фактичний обсяг вибірки для критерію n_{eff} може бути меншим за початкову кількість пар.

2.14. Розмір ефекту

Як і в параметричному аналізі, бажано оцінювати не лише статистичну значущість, а й силу ефекту.

Для Mann–Whitney за абсолютною величиною можна використати:

$$|r_{rb}| = 1 - \frac{2U}{n_1 n_2}.$$

Для Wilcoxon:

$$r_{rb} = \frac{T_+ - T_-}{T_+ + T_-}.$$

Знак показує напрям ефекту, абсолютна величина — його силу.

2.15. Медіана та IQR

Для непараметричних даних поряд із тестом доцільно подавати медіану, Q_1 , Q_3 та IQR. Це часто інформативніше, ніж повідомлення лише середнього.

2.16. Загальний алгоритм вибору критерію

незалежні групи → Mann–Whitney;	парні дані → Wilcoxon signed-rank
---------------------------------	-----------------------------------

Ключове питання перед розрахунком: мої вибірки незалежні чи парні?

3. Порядок виконання роботи

Частина А. Критерій Манна–Уїтні

1. Записати дві незалежні вибірки органолептичних оцінок.
2. Визначити n_1 і n_2 .
3. Обчислити медіани та квантілі.
4. Побудувати boxplot.
5. Сформулювати H_0 та H_1 .
6. Об'єднати дві вибірки та відсортувати дані.
7. Присвоїти ранги, використовуючи середні ранги для ties.

8. Знайти R_1 і R_2 .
9. Обчислити U_1 , U_2 та U .
10. Порівняти з критичним значенням або визначити p-value.
11. Обчислити ранговий розмір ефекту.
12. Перевірити результат у Microsoft Excel / Real Statistics / Jamovi.
13. Сформулювати статистичний і предметний висновки.

Частина Б. Критерій Уілкоксона

1. Записати дві парні вибірки.
2. Для кожної пари обчислити d_i .
3. Вилучити пари з $d_i=0$.
4. Обчислити $|d_i|$.
5. Ранжувати абсолютні різниці.
6. Урахувати ties через середні ранги.
7. Повернути знаки.
8. Знайти T_+ та T_- .
9. Визначити T .
10. Визначити p-value або критичне значення.
11. Обчислити ранговий effect size.
12. Виконати програмну перевірку.
13. Сформулювати статистичний і предметний висновки.

4. Приклад виконання. Демонстраційний варіант 0

Частина А. Критерій Манна–Уїтні

Досліджується загальна органолептична прийнятність двох варіантів десертного продукту. Продукт А оцінює одна група з 10 дегустаторів, продукт В — інша незалежна група з 10 дегустаторів. Використовується 9-бальна шкала.

№	Продукт А	Продукт В
1	7	5
2	8	6
3	6	6
4	9	7
5	7	5
6	8	6
7	7	7
8	6	5
9	8	6
10	7	6

Приймаємо $\alpha = 0.05$.

4.1. Описова характеристика

$$n_A = n_B = 10.$$

$$Me_A = 7, \quad Me_B = 6.$$

На описовому рівні продукт А має вищі оцінки, але потрібно визначити статистичну значущість.

4.2. Формулювання гіпотез

H_0 : систематичної різниці між розподілами оцінок А і В немає.

H_1 : розподіли органолептичних оцінок А і В відрізняються.

4.3. Об'єднання та ранжування

В об'єднаній вибірці є значення 5, 6, 7, 8, 9.

Значення 5 зустрічається 3 рази і займає ранги 1,2,3. Середній ранг:

$$\frac{1 + 2 + 3}{3} = 2.$$

Тому $5 \rightarrow 2$.

Значення 6 зустрічається 7 разів і займає позиції 4–10. Середній ранг:

$$\frac{4 + 10}{2} = 7.$$

Тому $6 \rightarrow 7$.

Значення 7 зустрічається 6 разів і займає позиції 11–16:

$$\frac{11 + 16}{2} = 13.5.$$

Тому $7 \rightarrow 13.5$.

Значення 8 займає ранги 17,18,19, середній ранг 18. Значення 9 має ранг 20.

Значення	Середній ранг
5	2.0
6	7.0
7	13.5
8	18.0
9	20.0

4.4. Сума рангів продукту A

Оцінки A: 7, 8, 6, 9, 7, 8, 7, 6, 8, 7. Відповідні ранги: 13.5, 18, 7, 20, 13.5, 18, 13.5, 7, 18, 13.5.

$$R_A = 142.$$

4.5. Сума рангів продукту B

Сума всіх рангів від 1 до 20:

$$1 + 2 + \dots + 20 = \frac{20 \cdot 21}{2} = 210.$$

Тому:

$$R_B = 210 - 142 = 68.$$

Перевірка: $R_A + R_B = 210$.

4.6. Обчислення U

$$U_A = R_A - \frac{10 \cdot 11}{2} = 142 - 55 = 87.$$

$$U_B = R_B - \frac{10 \cdot 11}{2} = 68 - 55 = 13.$$

Перевірка:

$$U_A + U_B = 87 + 13 = 100 = n_A n_B.$$

Розрахункова статистика:

$$\boxed{U = 13}.$$

4.7. Статистичне рішення

Для $n_A = n_B = 10$ та двостороннього $\alpha=0.05$ критичне значення $U_{crit} = 23$. Оскільки:

$$13 < 23,$$

Но відхиляємо. Програмний розрахунок із поправкою на ties і неперервність дає приблизно:

$$p \approx 0.004.$$

Відмінність статистично значуща.

4.8. Розмір ефекту

$$|r_{rb}| = 1 - \frac{2U}{n_A n_B} = 1 - \frac{26}{100} = 0.74.$$

Це свідчить про виражене розділення двох груп оцінок.

4.9. Статистичний і предметний висновки

Статистичний висновок: $U = 13$, $p \approx 0.004$, $r_{rb} \approx 0.74$. Продукт А характеризується вищими рангами оцінок.

Предметний висновок: у досліджених незалежних групах дегустаторів продукт А отримує вищі оцінки загальної прийнятності. Однак статистичний результат не замінює оцінки рецептури, собівартості, харчової цінності та інших характеристик.

Частина Б. Критерій знакових рангів Уїлкоксона

Дванадцять дегустаторів оцінили продукт до та після зміни рецептури за 9-бальною шкалою.

Дегустатор	До зміни	Після зміни
1	6	7
2	7	9
3	5	6
4	6	6
5	7	8
6	6	5
7	5	7
8	7	8
9	6	6
10	5	6
11	7	8
12	6	8

4.10. Формулювання гіпотез

H_0 : систематичного зсуву парних оцінок немає.

H_1 : парні оцінки після зміни рецептури систематично відрізняються.

4.11. Обчислення різниць

Використовуємо $d_i = x_{i,ext\text{після}} - x_{i,ext\text{до}}$.

№	До	Після	d_i
1	6	7	+1
2	7	9	+2
3	5	6	+1
4	6	6	0
5	7	8	+1
6	6	5	-1
7	5	7	+2
8	7	8	+1
9	6	6	0
10	5	6	+1

№	До	Після	d_i
11	7	8	+1
12	6	8	+2

Дві різниці дорівнюють 0 і вилучаються з ранжування:

$$n_{\text{eff}} = 10.$$

4.12. Ранжування абсолютних різниць

Серед ненульових різниць $|d| = 1$ зустрічається 7 разів і займає ранги 1–7. Середній ранг:

$$\frac{1 + 7}{2} = 4.$$

Значення $|d| = 2$ зустрічається 3 рази і займає ранги 8,9,10. Середній ранг:

$$\frac{8 + 10}{2} = 9.$$

4.13. Таблиця знакових рангів

№	d_i	d_i	Ранг	Знаковий ранг
1	+1	1	4	+4
2	+2	2	9	+9
3	+1	1	4	+4
5	+1	1	4	+4
6	-1	1	4	-4
7	+2	2	9	+9
8	+1	1	4	+4
10	+1	1	4	+4
11	+1	1	4	+4
12	+2	2	9	+9

4.14. Суми позитивних та негативних рангів

$$T_+ = 51.$$

$$T_- = 4.$$

Перевірка:

$$T_+ + T_- = 55 = \frac{10 \cdot 11}{2}.$$

Статистика критерію:

$$T = \min(51, 4) = 4.$$

4.15. Статистичне рішення

Програмний двосторонній розрахунок з урахуванням ties та нульових пар дає $p < 0.05$; залежно від використаної реалізації поправок приблизно 0.01–0.02. За нормальним наближенням з поправками можна отримати близько $p \approx 0.015$.

Отже, H_0 відхиляємо: після зміни рецептури органолептичні оцінки систематично змінилися.

4.16. Розмір ефекту

$$r_{rb} = \frac{T_+ - T_-}{T_+ + T_-} = \frac{51 - 4}{55} = \frac{47}{55} \approx 0.855.$$

Додатний знак показує переважання позитивних змін.

4.17. Описова характеристика пар

$$Me_{\text{до}} = 6, \quad Me_{\text{після}} = 7, \quad Me_d = 1.$$

4.18. Статистичний і предметний висновки

Статистичний висновок: $T = 4$, $p < 0.05$, $r_{rb} \approx 0.855$. Позитивні ранги суттєво переважають негативні.

Предметний висновок: після зміни рецептури більшість дегустаторів підвищила оцінку продукту; медіана збільшилася приблизно з 6 до 7 балів.

5. Виконання та перевірка в Microsoft Excel

Стандартний Excel не містить окремих базових команд Mann–Whitney та Wilcoxon signed-rank у ToolPak. Тому використовується схема: ручне ранжування → формули Excel → Real Statistics або Jamovi → зіставлення.

Для ранжування об'єднаних даних: `=RANK.AVG(value,all_values,1)`. Для суми рангів певної групи: `=SUMIF(group_range,"A",rank_range)`.

Для Wilcoxon створити стовпці: До, Після, d , $|d|$, Ранг, Знаковий ранг. Нульові різниці не ранжуються.

Real Statistics може використовувати функції типу MANN, MWTEST для Mann–Whitney та SRANK, SRTEST для Wilcoxon signed-rank. У Jamovi обирається непараметричний тест для незалежних або парних вибірок.

6. Вимоги до оформлення звіту

Для частини Mann–Whitney необхідно подати медіани, об'єднану рангову таблицю, R_1 , R_2 , U_1 , U_2 , U , критичне значення або p -value, effect size, статистичний та предметний висновки. Для Wilcoxon — парні спостереження, d_i , $|d_i|$, обробку нульових різниць, ранги, T_+ , T_- , T , p -value, effect size та два висновки.

7. Контрольні питання

1. У яких ситуаціях доцільно використовувати непараметричні критерії замість параметричних?
2. Чому органолептичні бали часто доцільно аналізувати ранговими методами?
3. Що таке ранг і як присвоюються ранги однаковим значенням?
4. У чому принципова різниця між критеріями Манна–Уїтні та Уїлкоксона?
5. Чому Mann–Whitney не можна використовувати для оцінок тих самих дегустаторів до і після зміни рецептури?
6. Чи завжди Mann–Whitney можна називати критерієм рівності медіан?
7. Що означають U_1 і U_2 і чому $U_1 + U_2 = n_1 n_2$?
8. Що відбувається з парами, для яких у Wilcoxon $d_i = 0$?
9. Чому у Wilcoxon ранжують абсолютні значення різниць, а потім повертають знаки?
10. Що показують T_+ і T_- ?
11. Навіщо поряд із p -value подавати медіани та розмір ефекту?
12. Чи означає статистично значуща перевага органолептичних оцінок, що нову рецептуру слід автоматично впроваджувати у виробництво?

Практична робота №6

Дисперсійний (ANOVA) та кореляційний аналіз технологічних даних

1. Мета роботи

Набути практичних навичок порівняння декількох технологічних режимів за допомогою однофакторного дисперсійного аналізу та дослідження статистичного зв'язку між технологічним фактором і результативним показником.

- перевіряти гіпотезу про рівність середніх трьох і більше незалежних технологічних груп за допомогою однофакторного ANOVA;
- розкласти загальну варіацію результатів на міжгрупову та внутрішньогрупову складові й обчислювати F-статистику;
- оцінювати величину групового ефекту за коефіцієнтом η^2 та інтерпретувати результати post-hoc порівнянь;
- будувати діаграму розсіювання та визначати коефіцієнти кореляції Пірсона і Спірмена засобами статистичного ПЗ;
- розрізняти висновки ANOVA і кореляційного аналізу та формулювати предметний технологічний висновок.

2. Короткі теоретичні відомості

2.1. Порівняння декількох технологічних режимів

У технологічному експерименті часто досліджують три або більше варіантів: температурні режими, рецептури, концентрації добавки, способи оброблення тощо.

Нехай досліджується k незалежних груп із генеральними середніми $\mu_1, \mu_2, \dots, \mu_k$. Перевіряється:

$$H_0: \mu_1 = \mu_2 = \dots = \mu_k$$

проти альтернативної гіпотези:

$$H_1: \text{принаймні одне генеральне середнє відрізняється}.$$

2.2. Чому не слід використовувати багато t-критеріїв

Якщо є чотири групи, число всіх попарних порівнянь:

$$C_4^2 = 6.$$

Якщо кожну пару перевіряти окремим t-тестом при $\alpha=0.05$, сумарна ймовірність хоча б однієї помилки першого роду зростає. Тому спочатку виконують один загальний тест — однофакторний ANOVA.

2.3. Основна ідея ANOVA

ANOVA порівнює два джерела варіативності: **міжгрупову** — відмінності між середніми технологічних режимів, та **внутрішньогрупову** — випадкове розсіювання всередині кожної групи.

Якщо вплив фактора відсутній, міжгрупова варіативність повинна бути співмірною з внутрішньогруповою. Якщо групові середні справді відрізняються, міжгрупова варіативність стає значно більшою.

2.4. Позначення

Нехай x_{ij} — j -те спостереження в i -й групі, n_i — обсяг i -ї групи, k — кількість груп.

$$N = \sum_{i=1}^k n_i.$$

Середнє i -ї групи:

$$\bar{x}_i = \frac{1}{n_i} \sum_{j=1}^{n_i} x_{ij}.$$

Загальне середнє:

$$\bar{x} = \frac{1}{N} \sum_{i=1}^k \sum_{j=1}^{n_i} x_{ij}.$$

2.5. Розкладання загальної варіації

Загальна сума квадратів:

$$SS_T = \sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x})^2.$$

Міжгрупова сума квадратів:

$$SS_B = \sum_{i=1}^k n_i (\bar{x}_i - \bar{x})^2.$$

Внутрішньогрупова сума квадратів:

$$SS_W = \sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x}_i)^2.$$

Виконується:

$$SS_T = SS_B + SS_W.$$

2.6. Ступені свободи

$$df_B = k - 1.$$

$$df_W = N - k.$$

$$df_T = N - 1.$$

І справедливо:

$$df_T = df_B + df_W.$$

2.7. Середні квадрати та F-критерій

$$MS_B = \frac{SS_B}{df_B}.$$

$$MS_W = \frac{SS_W}{df_W}.$$

Статистика ANOVA:

$$F = \frac{MS_B}{MS_W}.$$

Якщо $F_{emp} > F_{crit}$ або $p < \alpha$, нульову гіпотезу про рівність усіх середніх відхиляють.

2.8. ANOVA-таблиця

Джерело варіації	SS	df	MS	F	p-value
Між групами	SS_B	$k - 1$	MS_B	F	p
Усередині груп	SS_W	$N - k$	MS_W	—	—
Загалом	SS_T	$N - 1$	—	—	—

2.9. Умови застосування однофакторного ANOVA

- спостереження незалежні;
- результативна ознака кількісна;
- розподіли у групах або залишки приблизно нормальні;
- генеральні дисперсії груп приблизно однакові.

Перед ANOVA доцільно аналізувати boxplot, виконати Shapiro–Wilk та перевірку однорідності дисперсій, наприклад Levene test. У ПР №6 ці процедури виконуються програмно, оскільки їх логіку вже було опрацьовано у ПР №4.

2.10. Що означає статистично значущий ANOVA

Якщо $p < 0.05$, можна зробити висновок: не всі генеральні середні однакові. Але ANOVA **не** показує, які саме групи відрізняються. Для цього після значущого ANOVA застосовують post-hoc порівняння.

2.11. Post-hoc Tukey HSD

У цій роботі використовується Tukey HSD. Сам алгоритм вручну не розгортається — студент отримує таблицю попарних порівнянь у Real Statistics або Jamovi та повинен правильно її прочитати й інтерпретувати.

2.12. Розмір ефекту ANOVA

p -value відповідає на питання про статистичну значущість, але не показує, яка частка варіативності пов'язана з фактором. Для цього використаємо:

$$\eta^2 = \frac{SS_B}{SS_T}.$$

Наприклад, $\eta^2 = 0.70$ означає, що приблизно 70% вибіркової варіативності результату асоційовано з належністю до досліджуваних груп.

2.13. Кореляційний аналіз як інший погляд на ті самі дані

Якщо рівні технологічного фактора мають числовий зміст — наприклад 70, 75, 80, 85 °C — можна поставити інше запитання: чи спостерігається статистична тенденція зміни результату зі збільшенням температури?

ANOVA розглядає температуру як групувальний фактор і запитує: «Чи відрізняються середні для різних режимів?»

Кореляція розглядає температуру як кількісну змінну і запитує: «Чи змінюється результат систематично зі збільшенням температури?»

Це принципово різні статистичні питання.

2.14. Коефіцієнт кореляції Пірсона

$$r = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum (x_i - \bar{x})^2 \sum (y_i - \bar{y})^2}}.$$

Має місце $-1 \leq r \leq 1$. Додатний r означає загальну тенденцію спільного зростання, від'ємний — протилежних змін. Чим ближче $|r|$ до 1, тим сильніший саме **лінійний** зв'язок.

2.15. Коефіцієнт Спірмена

Spearman correlation ґрунтується на рангах. Він доцільний для порядкових даних, монотонного, але не обов'язково лінійного зв'язку, ненормальних даних або ситуацій із викидами. У ПР №6 r_s обчислюється програмно, щоб не дублювати ручне ранжування з ПР №5.

2.16. Перевірка значущості кореляції

Для Pearson перевіряють:

$$H_0: \rho = 0, \quad H_1: \rho \neq 0.$$

У програмному пакеті отримують r і p -value. Якщо $p < 0.05$, лінійний кореляційний зв'язок статистично значущий.

2.17. Діаграма розсіювання

Перед інтерпретацією коефіцієнта кореляції обов'язково будують scatterplot. Він дозволяє побачити напрям, приблизну силу, нелінійність та викиди. Одне число r без графіка може бути оманливим.

2.18. Кореляція не означає причинності

Навіть велике і статистично значуще $|r|$ не доводить причинний вплив X на Y . Причинне трактування залежить від дизайну експерименту, рандомізації, контролю сторонніх факторів і предметного механізму.

кореляція $\neq$ причинність.

2.19. Загальний алгоритм роботи

Основна ручна частина:

технологічні режими $\rightarrow$ ANOVA $\rightarrow F \rightarrow p \rightarrow \eta^2 \rightarrow$ Tukey

Коротке програмне продовження:

числові рівні фактора $\rightarrow$ scatterplot $\rightarrow r_p, r_s \rightarrow$ інтерпретація

3. Порядок виконання роботи

Частина А. Основний розрахунок ANOVA

1. Записати результати для декількох технологічних режимів.
2. Визначити k , n_i та N .
3. Для кожної групи обчислити середнє, дисперсію та стандартне відхилення.
4. Побудувати boxplot.
5. За допомогою статистичного ПЗ перевірити нормальність і однорідність дисперсій.
6. Сформулювати H_0 : $\mu_1 = \mu_2 = \dots = \mu_k$.
7. Обчислити загальне середнє.
8. Обчислити SS_B .
9. Обчислити SS_W .
10. Визначити $SS_T = SS_B + SS_W$.
11. Визначити ступені свободи.
12. Обчислити MS_B і MS_W .
13. Обчислити F_{emp} .
14. Визначити F_{crit} та p -value.
15. Прийняти статистичне рішення.
16. Обчислити η^2 .
17. Якщо ANOVA значущий, виконати Tukey HSD у Real Statistics або Jamovi.
18. Сформулювати статистичний та предметний висновки.

Частина Б. Скорочений кореляційний аналіз

1. Для кожного спостереження записати числове значення технологічного фактора X .
2. Побудувати scatterplot $X \rightarrow Y$.
3. За допомогою Excel визначити Pearson r_P .
4. За допомогою Real Statistics або Jamovi визначити Spearman r_S .
5. Визначити відповідні p -value.
6. Порівняти Pearson і Spearman.
7. Пояснити, чому висновок кореляційного аналізу не є тотожним висновку ANOVA.

4. Приклад виконання. Демонстраційний варіант 0

Досліджується вплив температури технологічного оброблення на вихід готового продукту. Використано чотири температурні режими: 70, 75, 80 та 85 °C. Для кожного режиму проведено по шість незалежних експериментів.

Дослід	70 °C	75 °C	80 °C	85 °C
1	91.8	93.0	94.2	92.7
2	92.2	93.4	94.5	93.0
3	91.5	92.8	94.1	92.9
4	92.0	93.1	94.4	93.1
5	91.7	93.3	94.3	92.8
6	92.1	92.9	94.6	93.0

Результативна ознака Y — вихід готового продукту, %. Приймаємо $\alpha = 0.05$.

4.1. Формулювання гіпотез

$$H_0: \mu_{70} = \mu_{75} = \mu_{80} = \mu_{85}.$$

H_1 : принаймні одне середнє відрізняється.

4.2. Описова статистика

Температура	Mean, %	Variance
70 °C	91.8833	0.06967
75 °C	93.0833	0.05367
80 °C	94.3500	0.03500
85 °C	92.9167	0.02167

Уже на описовому рівні видно, що максимальне середнє має режим 80 °C.

4.3. Перевірка передумов

За Shapiro–Wilk для кожної групи програмно отримано p -value більше 0.05. За критерієм Левена:

$$p \approx 0.399 > 0.05.$$

Отже, дані не дають підстав відхиляти припущення про нормальність та однорідність дисперсій. Застосовуємо класичний однофакторний ANOVA.

4.4. Загальне середнє

$$N = 24.$$

$$\bar{x} = 93.0583\%.$$

4.5. Міжгрупова сума квадратів

$$SS_B = 6(91.8833 - 93.0583)^2 + 6(93.0833 - 93.0583)^2 + 6(94.3500 - 93.0583)^2 + 6(92.9167 - 93.0583)^2.$$

Після обчислень:

$$SS_B \approx 18.4183.$$

4.6. Внутрішньогрупова сума квадратів

Оскільки $SS_W = \sum (n_i - 1)s_i^2$ і в кожній групі $n=6$:

$$SS_W = 5(0.06967) + 5(0.05367) + 5(0.03500) + 5(0.02167).$$

$$SS_W \approx 0.9000.$$

4.7. Загальна сума квадратів

$$SS_T = SS_B + SS_W = 18.4183 + 0.9000 = 19.3183.$$

4.8. Ступені свободи

$$df_B = 4 - 1 = 3.$$

$$df_W = 24 - 4 = 20.$$

$$df_T = 24 - 1 = 23.$$

Перевірка: $3 + 20 = 23$.

4.9. Середні квадрати

$$MS_B = \frac{18.4183}{3} = 6.1394.$$

$$MS_W = \frac{0.9000}{20} = 0.0450.$$

4.10. F-статистика

$$F = \frac{6.1394}{0.0450} \approx 136.43.$$

Для $df_1 = 3$, $df_2 = 20$, $\alpha=0.05$:

$$F_{\text{crit}} \approx 3.098.$$

Оскільки:

$$136.43 > 3.098,$$

Но відхиляємо. Програмний результат:

$$p < 0.0001.$$

4.11. ANOVA-таблиця

Джерело	SS	df	MS	F	p
Між групами	18.4183	3	6.1394	136.43	<0.0001
Усередині груп	0.9000	20	0.0450	—	—
Загалом	19.3183	23	—	—	—

4.12. Розмір ефекту

$$\eta^2 = \frac{18.4183}{19.3183} \approx 0.953.$$

У межах цього експерименту приблизно 95.3% спостережуваної варіативності виходу продукту пов'язано з розподілом спостережень між досліджуваними температурними режимами.

4.13. Post-hoc Tukey HSD

Оскільки ANOVA статистично значущий, у Real Statistics або Jamovi виконуємо Tukey HSD. Отримуємо:

Пара режимів	Висновок
70–75 °С	статистично значуща різниця
70–80 °С	статистично значуща різниця
70–85 °С	статистично значуща різниця
75–80 °С	статистично значуща різниця
75–85 °С	статистично значущої різниці немає
80–85 °С	статистично значуща різниця

Режим 80 °С має статистично вищий середній вихід продукту порівняно з іншими досліджуваними режимами.

4.14. Предметний висновок ANOVA

Температурний режим статистично пов'язаний із виходом готового продукту. Найбільший середній вихід — 94.35% — отримано при 80 °С. Підвищення температури до 85 °С не дає подальшого зростання, а навпаки зменшує середній вихід до 92.92%. Отже, вплив температури не має вигляду простого безперервного зростання.

4.15. Кореляційний аналіз тих самих даних

Тепер для кожного з 24 спостережень температура розглядається як числова змінна X , а вихід — як Y . Scatterplot показує збільшення виходу від 70 до 80 °С, але зменшення при 85 °С. Тому зв'язок не виглядає строго лінійним.

Pearson correlation

За всіма 24 спостереженнями Microsoft Excel дає:

$$r_P \approx 0.5442.$$

Для перевірки $H_0: \rho = 0$ отримуємо приблизно:

$$p_P \approx 0.0060 < 0.05.$$

Отже, загальна лінійна тенденція статистично значуща, але її сила лише помірна.

Spearman correlation

За програмним розрахунком:

$$r_S \approx 0.5015.$$

$$p_S \approx 0.0125 < 0.05.$$

Ранговий аналіз також показує помірну додатну тенденцію.

4.16. Чому ANOVA дуже значущий, а кореляція лише помірна

Це не суперечність. ANOVA запитує: «Чи відрізняються середні чотирьох температурних режимів?» — і відповідь: так, дуже суттєво. Pearson correlation запитує: «Чи існує приблизно лінійне зростання або спадання виходу зі збільшенням температури?» — і відповідь: лише помірна додатна тенденція, оскільки після 80 °С вихід зменшується.

$$\text{сильний груповий ефект} \Rightarrow \text{сильна лінійна кореляція}.$$

4.17. Загальний статистичний висновок

Однофакторний ANOVA показав статистично значущі відмінності між режимами: $F(3,20) = 136.43$, $p < 0.0001$, $\eta^2 = 0.953$. Post-hoc аналіз виділяє 80 °С як режим із найвищим середнім виходом. Водночас $r_P = 0.544$ і $r_S = 0.502$ показують лише помірну тенденцію, оскільки залежність не є строго лінійною або монотонною.

4.18. Загальний технологічний висновок

Температура технологічного оброблення суттєво впливає на вихід продукту, але принцип «чим вища температура — тим більший вихід» у дослідженому діапазоні не виконується. Максимальний результат спостерігається при 80 °С, що вказує на можливу наявність оптимального діапазону температур.

5. Виконання та перевірка в Microsoft Excel

Для ANOVA дані розміщують у чотирьох стовпцях і виконують: **Data → Data Analysis → ANOVA: Single Factor**. Excel виводить SS, df, MS, F, P-value та F crit.

Розмір ефекту η^2 обчислюється окремо: $=SS_between/SS_total$. Tukey HSD виконується у Real Statistics або Jamovi.

Для кореляції ті самі дані перетворюють на два стовпці «Температура» і «Вихід» та будують Scatter.

Процедура	Excel
Pearson r	=CORREL(X_range,Y_range) або =PEARSON(X_range,Y_range)
t для значущості r	=r*SQRT((n-2)/(1-r^2))
p-value	=T.DIST.2T(ABS(t),n-2)

Spearman у цій роботі визначається програмно. За бажанням у Excel можна ранжувати обидва стовпці функцією RANK.AVG і обчислити CORREL між рангами.

6. Вимоги до оформлення звіту

Основна ANOVA-частина повинна містити вихідні дані, описову статистику груп, boxplot, результати перевірки передумов, H_0/H_1 , загальне середнє, SS_B, SS_W, SS_T, ступені свободи, MS_B, MS_W, F_emp, F_crit, p-value, ANOVA-таблицю, η^2 , результати Tukey HSD, статистичний і технологічний висновки.

Скорочена кореляційна частина містить таблицю X,Y, scatterplot, r_P та p_P, r_S та p_S, коротке порівняння Pearson і Spearman та пояснення відмінності між висновками ANOVA і кореляції.

7. Контрольні питання

1. Яку нульову гіпотезу перевіряє однофакторний ANOVA?
2. Чому при порівнянні чотирьох груп небажано замінювати ANOVA шістьма окремими t-тестами?
3. Що характеризують міжгрупова та внутрішньогрупова суми квадратів?
4. Чому статистика ANOVA визначається як $F=MS_B/MS_W$?
5. Які основні передумови класичного однофакторного ANOVA?
6. Чому статистично значущий ANOVA не показує, які саме групи відрізняються? Для чого потрібний Tukey HSD?
7. Що показує коефіцієнт η^2 ?
8. У чому принципова відмінність між питанням ANOVA і питанням кореляційного аналізу?
9. Чому перед інтерпретацією коефіцієнта кореляції необхідно побудувати scatterplot?
10. У чому відмінність між Pearson і Spearman correlation?
11. Як можлива ситуація, коли ANOVA показує дуже сильні відмінності між групами, а коефіцієнт лінійної кореляції є лише помірним?
12. Чому статистично значуща кореляція не є автоматичним доказом причинного зв'язку?

Практична робота №7

Побудова та оцінювання регресійної моделі: прогнозування терміну зберігання продукту

1. Мета роботи

Набути практичних навичок побудови, статистичного оцінювання та використання парної лінійної регресійної моделі для аналізу й прогнозування технологічних показників.

- будувати діаграму розсіювання та обґрунтовувати доцільність використання лінійної регресійної моделі;
- визначати параметри рівняння парної лінійної регресії методом найменших квадратів;
- оцінювати якість і статистичну значущість регресійної моделі за R^2 , t - та F -критеріями;
- обчислювати прогнозні значення та аналізувати залишки моделі;
- розрізнати довірчий інтервал для середнього відгуку та прогнозний інтервал для окремого нового спостереження;
- використовувати регресійну модель для оцінювання моменту досягнення заданого критичного рівня показника якості;
- перевіряти ручні розрахунки в Microsoft Excel та інтерпретувати результати у контексті харчових технологій.

2. Короткі теоретичні відомості

2.1. Завдання регресійного аналізу

Кореляційний аналіз відповідає переважно на запитання: чи існує статистичний зв'язок між двома величинами і наскільки він сильний? Регресійний аналіз робить наступний крок: як математично описати цей зв'язок і використати його для оцінювання або прогнозування результативної ознаки?

Приклади у харчових технологіях: вологість від тривалості сушіння; кислотність від часу ферментації; показник окиснення від часу зберігання; мікробіологічний показник від тривалості зберігання; вихід продукту від температури.

2.2. Факторна та результативна ознаки

Величину, яка використовується для пояснення або прогнозування іншої, позначають X і називають незалежною змінною, факторною ознакою або предиктором. Величину, поведінку якої моделюють, позначають Y і називають залежною змінною або відгуком.

Наприклад: X — тривалість зберігання, Y — мікробіологічний показник.

2.3. Парна лінійна регресія

Найпростіша модель:

$$Y = \beta_0 + \beta_1 X + \varepsilon.$$

За вибірковими даними отримують оцінене рівняння:

$$\hat{Y} = b_0 + b_1 X.$$

2.4. Зміст коефіцієнта b_1

Коефіцієнт b_1 показує, на скільки одиниць у середньому змінюється прогнозоване значення Y при збільшенні X на одну одиницю.

Якщо $b_1 > 0$, зв'язок додатний; якщо $b_1 < 0$, від'ємний.

2.5. Зміст коефіцієнта b_0

b_0 — прогнозоване значення результативної ознаки при $X = 0$. Його предметна інтерпретація має сенс лише тоді, коли $X=0$ є змістовним і належить до області дослідження.

2.6. Метод найменших квадратів

Для кожного спостереження (x_i, y_i) модель дає прогноз $\hat{y}_i = b_0 + b_1 x_i$. Різниця між фактичним і прогнозованим значенням:

$$e_i = y_i - \hat{y}_i$$

називається **залишком**. Метод найменших квадратів вибирає b_0, b_1 так, щоб мінімізувати суму квадратів залишків:

$$SSE = \sum_{i=1}^n (y_i - \hat{y}_i)^2.$$

2.7. Оцінки коефіцієнтів

Введемо:

$$S_{xx} = \sum_{i=1}^n (x_i - \bar{x})^2.$$

$$S_{xy} = \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y}).$$

Тоді коефіцієнт нахилу:

$$b_1 = \frac{S_{xy}}{S_{xx}}.$$

Вільний член:

$$b_0 = \bar{y} - b_1 \bar{x}.$$

2.8. Прогнозовані значення

$$\hat{y}_i = b_0 + b_1 x_i.$$

Для кожного спостереження залишок:

$$e_i = y_i - \hat{y}_i.$$

2.9. Розкладання варіації

Загальна варіація:

$$SST = \sum_{i=1}^n (y_i - \bar{y})^2.$$

Варіація, пояснена моделлю:

$$SSR = \sum_{i=1}^n (\hat{y}_i - \bar{y})^2.$$

Непояснена залишкова варіація:

$$SSE = \sum_{i=1}^n (y_i - \hat{y}_i)^2.$$

Виконується:

$$SST = SSR + SSE.$$

2.10. Коефіцієнт детермінації

$$R^2 = \frac{SSR}{SST} = 1 - \frac{SSE}{SST}.$$

Має місце $0 \leq R^2 \leq 1$. Наприклад, $R^2=0.90$ означає, що приблизно 90% вибіркової варіативності Y описується побудованою лінійною моделлю з фактором X .

Високе R^2 саме по собі не доводить причинного зв'язку, правильності математичної форми, відсутності систематичних залишків або надійності прогнозу за межами досліджуваного діапазону.

2.11. Зв'язок між кореляцією і регресією

Для парної лінійної регресії з вільним членом:

$$R^2 = r^2.$$

Кореляція характеризує силу і напрям лінійного зв'язку, а регресія задає конкретне рівняння прогнозу.

2.12. Оцінка залишкової дисперсії

$$s_e^2 = \frac{SSE}{n-2}.$$

Залишкове стандартне відхилення:

$$s_e = \sqrt{\frac{SSE}{n-2}}.$$

Воно характеризує типовий масштаб відхилення фактичних значень від регресійної прямої.

2.13. Перевірка значущості коефіцієнта регресії

Перевіряється:

$$H_0: \beta_1 = 0, \quad H_1: \beta_1 \neq 0.$$

Якщо $\beta_1 = 0$, лінійна залежність між X і Y відсутня.

Стандартна похибка коефіцієнта нахилу:

$$SE_{b_1} = \frac{s_e}{\sqrt{S_{xx}}}.$$

Статистика:

$$t = \frac{b_1}{SE_{b_1}}, \quad \nu = n - 2.$$

Якщо $p < \alpha$, коефіцієнт нахилу статистично значущий.

2.14. Довірчий інтервал для коефіцієнта нахилу

$$b_1 \pm t_{1-\alpha/2; n-2} SE_{b_1}.$$

Якщо 95% CI для b_1 не містить 0, це узгоджується зі статистично значущим лінійним зв'язком при $\alpha=0.05$.

2.15. F-критерій значущості регресійної моделі

Для парної регресії:

$$MS_R = \frac{SSR}{1} = SSR.$$

$$MS_E = \frac{SSE}{n-2}.$$

$$F = \frac{MS_R}{MS_E}.$$

Для парної лінійної регресії тест значущості нахилу та загальний F-тест еквівалентні:

$$F = t^2.$$

2.16. Аналіз залишків

Високе R^2 ще не гарантує адекватності моделі. Необхідно аналізувати залишки e_i . Для прийнятної лінійної моделі вони мають хаотично розташовуватися навколо нуля, не демонструвати чіткої криволінійної структури, систематичного збільшення розсіювання або очевидних аномальних точок.

2.17. Гомоскедастичність

Гомоскедастичність означає приблизно сталу дисперсію випадкової складової при різних значеннях X . Якщо розсіювання залишків різко змінюється зі зміною X , маємо гетероскедастичність, що може свідчити про невідповідність простої лінійної моделі.

2.18. Точковий прогноз

Для заданого $X = x_0$:

$$\hat{y}_0 = b_0 + b_1 x_0.$$

2.19. Довірчий інтервал для середнього відгуку

Якщо потрібно оцінити середнє значення Y для всіх об'єктів при заданому x_0 :

$$\hat{y}_0 \pm t_{\text{crit}} s_e \sqrt{\frac{1}{n} + \frac{(x_0 - \bar{x})^2}{S_{xx}}}.$$

2.20. Прогнозний інтервал для нового спостереження

Якщо потрібно спрогнозувати значення для одного нового об'єкта:

$$\hat{y}_0 \pm t_{\text{crit}} s_e \sqrt{1 + \frac{1}{n} + \frac{(x_0 - \bar{x})^2}{S_{xx}}}.$$

Prediction interval ширший за confidence interval, оскільки включає додаткову індивідуальну варіативність нового спостереження.

2.21. Чому інтервали розширюються далеко від середнього X

В обох формулах присутній доданок $(x_0 - \bar{x})^2$. Найточніше модель оцінює відгук поблизу $\bar{x}$, а невизначеність збільшується біля країв області експерименту.

2.22. Інтерполяція та екстраполяція

Інтерполяція — прогноз усередині дослідженого діапазону X . **Екстраполяція** — прогноз за його межами. Екстраполяція значно ризикованіша, оскільки невідомо, чи збережеться встановлена закономірність поза областю експерименту.

2.23. Оцінювання терміну досягнення критичного рівня

Нехай показник у часі описується моделлю $\hat{Y} = b_0 + b_1 t$. Якщо критичне значення становить Y_{crit} , момент його прогнозованого досягнення визначається з рівняння:

$$Y_{\text{crit}} = b_0 + b_1 t_{\text{crit}}.$$

Тому:

$$t_{\text{crit}} = \frac{Y_{\text{crit}} - b_0}{b_1}.$$

Це точкова модельна оцінка. Для нормативного встановлення терміну придатності потрібні належний дизайн, достатній обсяг даних, оцінювання невизначеності, перевірка моделі, нормативні вимоги та запас безпеки.

2.24. Коли лінійна модель непридатна

Ознаки можливої непридатності: явно криволінійний scatterplot, систематичний вигин залишків, сильна гетероскедастичність, впливові викиди, нелогічні прогнози або відома з предметних міркувань нелінійна природа процесу. У таких випадках можуть використовуватися поліноміальні, показникові, логарифмічні, степеневі або інші моделі.

2.25. Загальний алгоритм регресійного аналізу

постановка задачі → scatterplot → вибір форми моделі → $b_0, b_1 \rightarrow R^2$

значущість → залишки → прогноз → висновок

3. Порядок виконання роботи

1. Записати вихідні парні спостереження (x_i, y_i) .
2. Визначити зміст незалежної та залежної змінних.
3. Побудувати діаграму розсіювання.
4. На підставі графіка оцінити доцільність лінійної моделі.
5. Обчислити $\bar{x}$ та $\bar{y}$.
6. Обчислити S_{xx} і S_{xy} .
7. Визначити b_1 та b_0 .
8. Записати рівняння $\hat{Y} = b_0 + b_1 X$.
9. Для кожного спостереження знайти $\hat{y}_i$.
10. Обчислити залишки e_i .
11. Визначити SST, SSR, SSE.
12. Перевірити $SST = SSR + SSE$.
13. Обчислити R^2 .
14. Обчислити залишкове стандартне відхилення s_e .
15. Перевірити статистичну значущість b_1 .
16. Побудувати 95% CI для b_1 .
17. Перевірити загальну значущість моделі за F.
18. Побудувати графік залишків і оцінити адекватність лінійної моделі.
19. Для заданого x_0 виконати точковий прогноз.
20. Побудувати 95% CI середнього відгуку і 95% prediction interval.
21. За заданого критичного значення Y визначити модельну оцінку моменту його досягнення.
22. Повторити розрахунки у Microsoft Excel.
23. Зіставити результати та сформулювати статистичний і предметний висновки.

4. Приклад виконання. Демонстраційний варіант 0

Під час дослідження зберігання харчового продукту періодично визначали логарифм загальної кількості життєздатних мікроорганізмів. Для навчального прикладу приймаємо критичний рівень:

$$Y_{\text{crit}} = 6.0 \log_{10} \text{ КУО/г.}$$

Час зберігання t, діб	Y, log10 КУО/г
0	2.1
2	2.5
4	3.0
6	3.6
8	4.1
10	4.7
12	5.2
14	5.8
16	6.3

4.1. Первинний графічний аналіз

Scatterplot утворює виражену висхідну приблизно лінійну структуру. Тому для демонстраційного аналізу приймаємо модель:

$$Y = \beta_0 + \beta_1 t + \varepsilon.$$

4.2. Середні значення

$$n = 9.$$

Середній час:

$$\bar{x} = \frac{0 + 2 + 4 + \dots + 16}{9} = 8 \text{ діб.}$$

Середній мікробіологічний показник:

$$\bar{y} = \frac{2.1 + 2.5 + 3.0 + 3.6 + 4.1 + 4.7 + 5.2 + 5.8 + 6.3}{9} = 4.1444.$$

4.3. Обчислення S_{xx}

$$S_{xx} = \sum (x_i - \bar{x})^2 = 240.$$

4.4. Обчислення S_{xy}

$$S_{xy} = \sum (x_i - \bar{x})(y_i - \bar{y}) = 64.4.$$

4.5. Коефіцієнт нахилу

$$b_1 = \frac{64.4}{240} = 0.26833.$$

4.6. Вільний член

$$b_0 = 4.14444 - 0.26833 \cdot 8 = 1.99778.$$

4.7. Рівняння регресії

$$\hat{Y} = 1.9978 + 0.26833t.$$

Інтерпретація b_1 : збільшення терміну зберігання на одну добу асоціюється із середнім збільшенням досліджуваного мікробіологічного показника приблизно на 0.268 log10 КУО/г.

4.8. Прогнозовані значення і залишки

Для $t=0$:

$$\hat{Y} = 1.9978.$$

Залишок:

$$e = 2.1 - 1.9978 = 0.1022.$$

Аналогічно обчислюємо для інших спостережень:

t	Y	$\hat{Y}$	$e=Y-\hat{Y}$
0	2.1	1.9978	0.1022
2	2.5	2.5344	-0.0344
4	3.0	3.0711	-0.0711
6	3.6	3.6078	-0.0078
8	4.1	4.1444	-0.0444
10	4.7	4.6811	0.0189
12	5.2	5.2178	-0.0178
14	5.8	5.7544	0.0456
16	6.3	6.2911	0.0089

Сума залишків у моделі з вільним членом практично дорівнює нулю.

4.9. Загальна варіація

$$SST = \sum (y_i - \bar{y})^2 \approx 17.30222.$$

4.10. Пояснена варіація

$$SSR = \sum (\hat{y}_i - \bar{y})^2 \approx 17.28067.$$

4.11. Залишкова варіація

$$SSE = \sum (y_i - \hat{y}_i)^2 \approx 0.02156.$$

Перевірка:

$$17.28067 + 0.02156 \approx 17.30222.$$

4.12. Коефіцієнт детермінації

$$R^2 = \frac{17.28067}{17.30222} \approx 0.99875.$$

Тобто в межах цієї вибірки приблизно 99.88% варіативності Y описується його лінійною залежністю від часу зберігання. Це дуже високе значення, але його потрібно розглядати разом із залишками.

4.13. Коефіцієнт кореляції

Для парної регресії $R^2 = r^2$. Оскільки $b_1 > 0$:

$$r \approx +\sqrt{0.99875} \approx 0.99938.$$

Це відповідає дуже сильному додатному лінійному зв'язку.

4.14. Залишкове стандартне відхилення

$$s_e = \sqrt{\frac{0.02156}{9-2}} = \sqrt{\frac{0.02156}{7}} \approx 0.05549.$$

4.15. Стандартна похибка b_1

$$SE_{b_1} = \frac{0.05549}{\sqrt{240}} \approx 0.003582.$$

4.16. t-критерій для коефіцієнта нахилу

$$H_0: \beta_1 = 0, \quad H_1: \beta_1 \neq 0.$$

$$t = \frac{0.26833}{0.003582} \approx 74.91.$$

Ступені свободи: $df = 7$. Для $\alpha=0.05$:

$$t_{\text{crit}} = t_{0.975;7} \approx 2.365.$$

Оскільки $74.91 \gg 2.365$, H_0 відхиляємо. Програмно:

$$p < 0.0001.$$

Коефіцієнт нахилу статистично значущий.

4.17. 95% довірчий інтервал для b_1

$$b_1 \pm t_{\text{crit}} SE_{b_1} = 0.26833 \pm 2.3646 \cdot 0.003582.$$

$$95\% CI_{\beta_1} = [0.2599; 0.2768].$$

Інтервал не містить 0.

4.18. F-критерій значущості моделі

$$df_R = 1, \quad df_E = 7.$$

$$MS_R = SSR = 17.28067.$$

$$MS_E = \frac{0.02156}{7} \approx 0.003079.$$

$$F = \frac{17.28067}{0.003079} \approx 5611.76.$$

Програмно $p < 0.0001$, отже регресійна модель статистично значуща.

Перевірка еквівалентності:

$$74.91^2 \approx 5611.76.$$

4.19. ANOVA-таблиця регресії

Джерело	SS	df	MS	F	p
Регресія	17.28067	1	17.28067	5611.76	<0.0001
Залишок	0.02156	7	0.003079	—	—
Загалом	17.30222	8	—	—	—

4.20. Аналіз залишків

Залишки змінюються приблизно від -0.071 до 0.102 , мають як додатні, так і від'ємні значення і не демонструють очевидної систематичної тенденції. На графіку e_i проти t_i не повинно спостерігатися вираженої криволінійної структури. Для демонстраційних даних лінійна модель є прийнятним описом.

4.21. Прогноз на 15-ту добу

$$\hat{Y}_{15} = 1.9978 + 0.26833 \cdot 15.$$

$$\hat{Y}_{15} \approx 6.0228 \log_{10} \text{ КУО/Г}.$$

4.22. 95% довірчий інтервал середнього відгуку на 15-ту добу

$$\bar{x} = 8, \quad S_{xx} = 240, \quad n = 9.$$

$$SE_{\text{mean}} = 0.05549 \sqrt{\frac{1}{9} + \frac{(15-8)^2}{240}} \approx 0.03116.$$

Тоді:

$$6.0228 \pm 2.3646 \cdot 0.03116.$$

$$95\% CI_{\text{mean}} \approx [5.949; 6.096].$$

4.23. 95% прогнознiй iнтервал для нового спостереження

$$SE_{\text{pred}} = 0.05549 \sqrt{1 + \frac{1}{9} + \frac{(15 - 8)^2}{240}} \approx 0.06364.$$
$$6.0228 \pm 2.3646 \cdot 0.06364.$$

$95\% PI \approx [5.872; 6.173]$

Як i очiкувалося, prediction interval ширший за CI середнього вiдгуку.

4.24. Оцiнювання термiну досягнення критичного рiвня

У навчальнiй задачi $Y_{\text{crit}} = 6.0$. Пiдставляємо в модель:

$$6.0 = 1.9978 + 0.26833t.$$

Тодi:

$$t_{\text{crit}} = \frac{6.0 - 1.9978}{0.26833} \approx 14.92 \text{ доби.}$$

За побудованою моделлю критичний рiвень досягається приблизно на 15-ту добу.

4.25. Iнтерполяцiя чи екстраполяцiя

Експериментальнi данi отриманi для $0 \leq t \leq 16$. Оцiнка $t_{\text{crit}} \approx 14.92$ лежить всерединi цього дiапазону, тому це **iнтерполяцiйна** оцiнка, методично надiйнiша за прогноз за межами експерименту.

4.26. Статистичний висновок

Побудовано модель $\hat{Y} = 1.9978 + 0.26833t$. Вона має $R^2 = 0.99875$. Коефiцiєнт нахилу статистично значущий: $t(7) = 74.91$, $p < 0.0001$. Загальна модель також статистично значуща: $F(1,7) = 5611.76$, $p < 0.0001$. Аналiз залишкiв не виявляє очевидної систематичної структури.

4.27. Предметний технологiчний висновок

У дослідженому дiапазонi часу мiкробiологiчний показник виражено збiльшується пiд час зберiгання. За моделлю критичний рiвень $6.0 \log_{10} \text{ КУО/г}$ досягається приблизно через 14.9 доби. Це модельна статистична оцiнка, а не автоматично встановлений нормативний термiн придатностi. Практичне встановлення термiну придатностi повинно враховувати нормативи, невизначенiсть прогнозу, умови зберiгання, повторюванiсть результатiв, запас безпеки та iншi показники якостi.

5. Виконання та перевiрка в Microsoft Excel

Основний принцип: ручний розрахунок $\rightarrow$ Excel $\rightarrow$ аналiз виходу Regression $\rightarrow$ зiставлення $\rightarrow$ прогноз.

Показник	Excel
Нахил b_1	=SLOPE(B2:B10,A2:A10)
Вiльний член b_0	=INTERCEPT(B2:B10,A2:A10)
R^2	=RSQ(B2:B10,A2:A10)
Кореляцiя r	=CORREL(A2:A10,B2:B10)
Стандартна похибка регресiї	=STEYX(B2:B10,A2:A10)
Прогноз на 15-ту добу	=FORECAST.LINEAR(15,B2:B10,A2:A10)
Розширена статистика	=LINEST(B2:B10,A2:A10,TRUE,TRUE)

Повний аналiз: **Data $\rightarrow$ Data Analysis $\rightarrow$ Regression**. Доцiльно активувати Residuals, Residual Plots та Line Fit Plots. У вихiдних таблицях потрiбно вмити знайти R Square, Adjusted R Square, Standard Error, ANOVA, F, Significance F, коефiцiєнти, iх стандартнi похибки, t Stat, P-value та 95% CI.

5.1. Довірчий інтервал середнього прогнозу в Excel

Для $x_0 = 15$ використовують середнє $\bar{X}$, S_{xx} , s_e та критичне t . S_{xx} можна отримати функцією =DEVSQ(A2:A10).

Стандартна похибка середнього прогнозу:
=STEYX(B2:B10,A2:A10)*SQRT(1/COUNT(A2:A10)+(15-AVERAGE(A2:A10))^2/DEVSQ(A2:A10)).

5.2. Прогнозний інтервал

Стандартна похибка нового спостереження:
=STEYX(B2:B10,A2:A10)*SQRT(1+1/COUNT(A2:A10)+(15-AVERAGE(A2:A10))^2/DEVSQ(A2:A10)). Далі до точкового прогнозу додають/віднімають $t_{crit}SE_{pred}$.

5.3. Ознайомче розширення: множинна регресія

У реальному технологічному процесі відгук часто залежить від кількох факторів:

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k + \varepsilon.$$

У межах цієї практичної роботи ручний матричний розрахунок множинної регресії **не виконується**. За наявності часу студент може ознайомитися з нею через Data Analysis → Regression із кількома X-стовпцями, звертаючи увагу на R^2 , Adjusted R^2 , Significance F, коефіцієнти, p-value та залишки. Поняття мультиколінеарності і VIF залишаються на лекційному рівні.

6. Вимоги до оформлення звіту

Звіт повинен містити постановку технологічної задачі, вихідні дані, визначення X та Y, scatterplot, обґрунтування лінійної моделі, ручні $\bar{x}$, $\bar{y}$, S_{xx} , S_{xy} , b_1 , b_0 , рівняння регресії, таблицю прогнозованих значень і залишків, SST, SSR, SSE, R^2 , s_e , t-перевірку b_1 , 95% CI b_1 , ANOVA-таблицю регресії, F-критерій, графік залишків, точковий прогноз, CI середнього прогнозу, prediction interval, оцінку критичного моменту, результати Excel, таблицю зіставлення, статистичний та предметний висновки.

7. Контрольні питання

1. У чому полягає відмінність між кореляційним та регресійним аналізом?
2. Що означають коефіцієнти b_0 та b_1 у рівнянні $\hat{Y} = b_0 + b_1 X$?
3. У чому полягає принцип методу найменших квадратів?
4. Що таке залишок і чому аналіз залишків необхідний навіть при дуже високому R^2 ?
5. Що показує коефіцієнт детермінації R^2 і чого він не показує?
6. Чому для парної лінійної регресії виконується $R^2 = r^2$?
7. Яку гіпотезу перевіряє t-критерій для b_1 ? Що означав би випадок $\beta_1 = 0$?
8. Яку гіпотезу перевіряє F-критерій моделі і чому для парної регресії $F = t^2$?
9. У чому різниця між 95% CI середнього відгуку і 95% prediction interval нового спостереження? Чому другий ширший?
10. У чому різниця між інтерполяцією та екстраполяцією? Чому екстраполяція ризикованіша?
11. Як за регресійним рівнянням визначити момент досягнення критичного значення показника і чому отримане число не можна автоматично ототожнювати з нормативним терміном придатності?
12. Чи можна на підставі високого R^2 зробити висновок про причинний вплив X на Y?

Рекомендована література

1. Шелестовський Б. Г., Габрусєв Г. В., Габрусєва І. Ю. Вища математика: теорія ймовірностей та математична статистика : навч. посіб. Тернопіль : СМП «Тайп», 2023. 142 с.
2. Плічко А. М., Акбаш К. С., Луньова М. В. Математична статистика : навчальний посібник. Кропивницький : Видавець Лисенко С. В., 2024. 220 с.
3. Білущак Г. І. Аналітичні та чисельні методи дослідження. Статистичні методи аналізу даних : навчальний посібник для здобувачів вищої освіти. Львів : Растр-7, 2024. 315 с.
4. Дударєв І. М., Кузьмін О. В. Практикум з методології наукових досліджень : навчальний посібник. Одеса : Олді+, 2023. 278 с.
5. Михайлицька О. Р., Сливка Н. Б., Наговська В. О., Білик О. Я. Методологія наукових досліджень : навчальний посібник для здобувачів вищої освіти спеціальності 181 «Харчові технології». Львів : Простір-М, 2024. 258 с.
6. Горбачук В. М., Кушлик-Дивульська О. І. Теорія ймовірностей та математична статистика : підручник для здобувачів ступеня бакалавра за технічними та економічними спеціальностями. Київ : КПІ ім. Ігоря Сікорського, 2023. 351 с.
7. Бідюк П. І., Данилов В. Я., Жиров О. Л. Прикладна статистика : навчальний посібник. Київ : КПІ ім. Ігоря Сікорського, 2023. 186 с.
8. Костюк В. О., Мількін І. В., Славута О. І. Статистика : підручник. Харків : ХНУМГ ім. О. М. Бекетова, 2023. 204 с.
9. Григорків В. С., Вінничук О. Ю., Григорків М. В., Маханець Л. Л. Статистика: основи теорії та практикум : навчальний посібник. Чернівці : Чернівецький нац. ун-т ім. Юрія Федьковича : Рута, 2022. 304 с.
10. Герасименко С. С., Голубова Г. В., Потапова М. Ю., Червона С. П. Статистика : навчальний посібник / за ред. О. Г. Осауленка. Київ : НАСОА, 2022. 265 с.
11. Загальна теорія статистики : підручник / І. А. Дмитрієв, О. А. Дмитрієва, О. М. Гіржева, А. В. Непран, Н. О. Бірченко, А. А. Воронкова, Н. В. Чуйко ; за ред. А. В. Непрана, І. А. Дмитрієва. Харків : ПП Іванченка, 2022. 720 с.