

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
ХАРКІВСЬКИЙ НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ
МІСЬКОГО ГОСПОДАРСТВА імені О. М. БЕКЕТОВА

Л. П. Вороновська

МАТЕМАТИЧНІ МЕТОДИ В ПСИХОЛОГІЧНИХ
ДОСЛІДЖЕННЯХ

Модуль 1

КОНСПЕКТ ЛЕКЦІЙ

*(для здобувачів першого (бакалаврського) рівня вищої освіти
денної та заочної форм навчання зі спеціальності
053 – Психологія)*

Харків
ХНУМГ ім. О. М. Бекетова
2025

УДК 519.2(316.6)

Вороновська Л. П. Математичні методи в психологічних дослідженнях. Модуль 1 : конспект лекцій для здобувачів першого (бакалаврського) рівня вищої освіти денної та заочної форм навчання зі спеціальності 053 – Психологія / Л. П. Вороновська ; Харків. нац. ун-т міськ. госп-ва ім. О. М. Бекетова. – Харків : ХНУМГ ім. О. М. Бекетова, 2025. – 126 с.

Автор

канд. пед. наук Л. П. Вороновська

Рецензент

Л. Б. Коваленко, кандидат фізико-математичних наук, доцент, завідувачка кафедри вищої математики і математичного моделювання (Харківський національний університет міського господарства імені О. М. Бекетова)

Рекомендовано кафедрою вищої математики і математичного моделювання, протокол № 5 від 07 листопада 2024 року

Конспект лекцій складено з метою допомогти здобувачам вишів за напрямом підготовки з психології під час занять та виконання завдань з математичних методів в психологічних дослідженнях.

© Л. П. Вороновська, 2025

© ХНУМГ ім. О. М. Бекетова, 2025

ЗМІСТ

ВСТУП.....	5
ЛЕКЦІЯ 1 ОСНОВНІ ПОНЯТТЯ МАТЕМАТИЧНОЇ СТАТИСТИКИ.....	6
1.1 Змінні та їх вимірювання. Вимірювальні шкали.....	6
1.2 Поняття про залежність між змінними. Залежні та незалежні змінні.....	8
1.3 Основні символи змінних та операцій. Поняття масиву даних.....	9
ЛЕКЦІЯ 2 СТАТИСТИЧНЕ СПОСТЕРЕЖЕННЯ. ПОНЯТТЯ ПРО РЕПРЕЗЕНТАТИВНІСТЬ. ГРАФІЧНЕ ПРЕДСТАВЛЕННЯ ДАНИХ.....	11
2.1 Статистичне спостереження: види, способи, принципи та помилки.....	11
2.2 Генеральна сукупність та вибірка. Репрезентативність вибірki.....	13
2.3 Математичні показники вибірки	15
2.3.1 Емпіричний розподіл частот.....	15
2.3.2 Міри центральної тенденції.....	19
2.3.3 Міри мінливості.....	21
2.3.4 Характеристики асиметрії та ексцесу.....	22
2.4 Графічне зображення частотних розподілів	25
2.4.1 Полігон частот.....	25
2.4.2 Гістограма частот.....	26
2.4.3 Вплив графічних зображень на уявлення про сукупність.....	27
ЛЕКЦІЯ 3 ОСНОВНІ СТАТИСТИЧНІ РОЗПОДІЛИ.....	29
3.1 Різновиди форм емпіричних розподілів.....	29
3.2 Нормальні розподіли.....	31
3.3 Побудова кривої нормального розподілу за емпіричними даними.....	34
3.4 Перевірка нормальності розподілу результативної ознаки.....	36
ЛЕКЦІЯ 4 СТАНДАРТНИЙ НОРМАЛЬНИЙ РОЗПОДІЛ. ДЕЯКІ СПЕЦІАЛЬНІ РОЗПОДІЛИ.....	40
4.1 Стандартний нормальний розподіл.....	40
4.1.1 Стандартний нормальний розподіл.....	40

4.1.2 Біноміальні розподіли.....	43
4.1.3 Пуассонові розподіли.....	45
4.1.4 U-подібні розподіли.....	46
4.2 Деякі спеціальні розподіли.....	48
4.2.1 t-розподіл Стюдента.....	48
4.2.2 χ^2 - розподіл Пірсона.....	51
4.2.3 F-розподіл Фішера-Снедекора.....	53
ЛЕКЦІЯ 5 МІРИ ЗВ'ЯЗКУ МІЖ ОЗНАКАМИ.....	55
5.1 Загальні положення.....	55
5.2 Коефіцієнт рангової кореляції r_s Спірмена.....	63
5.3 Коефіцієнт кореляції Брава-Пірсона.....	70
5.4 Інтерпретація коефіцієнтів кореляції.....	71
ЛЕКЦІЯ 6 МЕТОДИ ПЕРЕВІРКИ СТАТИСТИЧНИХ ГІПОТЕЗ.....	73
6.1 t-критерій Стюдента.....	73
6.2 F-критерій Фішера (для порівняння дисперсій).....	78
6.3 Виявлення відмінностей у рівні досліджуваної ознаки. Q-критерій Розенбаума.....	80
6.4 T-критерій Вілкоксона.....	87
6.5 Виявлення відмінностей у розподілі ознаки. χ^2 -критерій Пірсона.....	92
ЛЕКЦІЯ 7 РЕГРЕСІЙНИЙ АНАЛІЗ.....	102
7.1 Задачі регресійного аналізу.....	102
7.2 Визначення коефіцієнтів регресії.....	103
7.2.1 Метод вибраних точок.....	105
7.2.2 Метод найменших квадратів.....	106
7.3 Обчислення похибки рівняння регресії.....	107
7.4 Види рівняння регресії.....	109
ЛЕКЦІЯ 8 ДИСПЕРСІЙНИЙ АНАЛІЗ.....	111
8.1 Поняття про факторні гіпотези.....	111
8.2 Типи факторних планів.....	112
8.3 Результати факторних експериментів.....	115
8.4 Двофакторний дисперсійний аналіз для незалежних вибірок.....	117
8.5 Проведення двофакторного дисперсійного аналізу. Критерій Фішера F	119
СПИСОК РЕКОМЕНДОВАНОЇ ЛІТЕРАТУРИ.....	125

ВСТУП

Конспект лекцій побудований за модульною технологією навчання згідно з робочою програмою курсу «Математичні методи в психологічних дослідженнях» для здобувачів першого (бакалаврського) рівня вищої освіти денної та заочної форм навчання зі спеціальності 053 – Психологія.

До конспекту увійшли лекції за темами «Аналіз статистичних даних», «Основні статистична розподіли», «Вибіркове дослідження», «Перевірка гіпотез», «Кореляційний аналіз», «Регресійний аналіз» та «Дисперсійний аналіз».

Доступне коротке подання теоретичного матеріалу супроводжується детальними ілюстраціями, великою кількістю прикладів для практичного закріплення вивченого.

Конспект лекцій може бути використаний на практичних заняттях для виконання домашніх завдань, а також для самостійного вивчення деяких розділів математичних методів в психологічних дослідженнях.

ЛЕКЦІЯ 1

ОСНОВНІ ПОНЯТТЯ МАТЕМАТИЧНОЇ СТАТИСТИКИ

1. Змінні та їх вимірювання. Вимірювальні шкали.
2. Поняття про залежність між змінними. Залежні та незалежні змінні.
3. Основні символи змінних та операцій. Поняття масиву даних.

1.1 Змінні та їх вимірювання. Вимірювальні шкали

До базових понять, які стосуються використання математичної статистики у психології, можна перш за все віднести поняття “змінні” та поняття “вимірювання”.

Що ж таке змінні? Фактично, *змінні* – це все, що можна вимірювати, контролювати або змінювати у дослідженнях. Змінною може бути рівень особистісної тривожності, інтелект, рівень агресії, мотивація досягнення, тип темпераменту тощо. Що буде змінною у Вашому дослідженні і скільки буде цих змінних, залежить від його теми, мети, гіпотези, методів, завдань та структури. Навіть у дослідженнях, які виконані на одну тему, змінні можуть бути різними, оскільки відрізнятимуться гіпотези чи використані для дослідження методики.

Вимірювання – це процес приписування чисел певним явищам (змінним) відповідно до певних правил. Таким чином, у процесі вимірювання ми маємо справу з трьома елементами:

- 1) числа;
- 2) явища (змінні);
- 3) правила приписування чисел явищам.

Найпоширенішою у психології є типологія шкал С. Стівенса, в основу якої покладено точність градування шкал та операції, які можна виконувати з числами. В межах цієї типології вирізняють такі типи вимірювальних шкал:

- шкала найменувань (номінальна),
- шкала порядку (рангова чи ординальна),
- шкала інтервалів (інтервальна),
- шкала відношень (пропорційна).

Дослідники Дж. Гласс та Дж. Стенлі наводять таблицю характеристик вимірювальних шкал С. Стівенса, яка має такий вигляд (табл. 1.1):

Таблиця 1.1

Шкала	Характеристика	Матем. операції	Приклад
1	2	3	4
Номінальна	Об'єкти класифікуються, а класи позначаються номерами. Таким чином, число тут слугує лише назвою певного класу, а тому нічого не говорить про властивості об'єкту, крім того, що він належить до певного класу. Частковим випадком шкали найменувань є дихотомічна шкала	Встановлення рівності	Колір очей, номери гравців спортивної команди, автомобільні чи інвентарні номери, типи темпераменту тощо
Рангова	Об'єкти розташовуються в порядку спадання чи зростання у них певної якості. При цьому кожній градації якості приписується свій порядковий номер (ранг). Фактично, об'єкти лише впорядковуються. Особливість шкали – однакові різниці між сусідніми рангами не означають однакової різниці між ступенями прояву виміряної якості	Встановлення відношень “більше-менше”	Нагороди за заслуги, військові ранги, місця на олімпіаді, ранжування досліджуваних за проявами індивідуальних рис

Продовження таблиці 1.1

1	2	3	4
Інтервальна	У межах цієї шкали існує одиниця вимірювання, з допомогою якої можна не лише впорядкувати об'єкти, але й приписати їм числа так, щоб однакові різниці між числами виражали однакові відмінності у проявах вимірюваної якості. Особливість шкали – довільний нуль, який не говорить про відсутність якості у об'єкта	Встановлення рівності інтервалів (різниць)	Календарний час, шкала температур за Цельсієм
Пропорційна	У межах цієї шкали числа мають такі ж властивості, як і в шкалі інтервалів, але, крім того, відношення чисел виражають кількісні відношення ступенів прояву явища. Особливість шкали – наявність абсолютного нуля, який означає відсутність явища	Встановлення рівності відношень	Зріст, вага, температура за Кельвіном, рівень інтелекту, мотивація тощо

1.2 Поняття про залежність між змінними.

Залежні та незалежні змінні

Першою базовою задачею є **пошук зв'язку** між явищами (змінними). Вважається, що дві чи більше змінних пов'язані між собою, коли їх значення розподілені узгоджено – коли систематично зміна однієї змінної в певному напрямку супроводжується зміною іншої змінної.

Друга задача (**аналіз впливу**) випливає із проблеми, які ставить перша задача – якщо між явищами існує зв'язок, то чи

означає це, що одне з них впливає на інше? А якщо так, то що на що впливає? Виявляється, методи пошуку зв'язку на ці питання не дають відповіді.

Виявляється, що для аналізу впливу одного явища на інше можна використати лише один дійсно надійний та валідний підхід – *проведення експерименту*. Суть експерименту полягає у тому, що дослідник систематично змінює одне з явищ, яке знаходиться у нього під контролем, і спостерігає, як буде змінюватися інше явище. У цьому випадку необхідно розрізнення двох типів змінних: *незалежних та залежних*.

Незалежна змінна – явище, яке знаходиться під контролем експериментатора, і яке він може змінювати відповідно до наперед визначеної експериментальної процедури.

Залежна змінна – явище, яке змінюється під впливом незалежної змінної.

Слід пам'ятати, що коли ми говоримо про *залежні змінні* – маємо на увазі *задачу пошуку зв'язку*, а коли говоримо про *залежну та незалежну змінні* – маємо на увазі *задачу аналізу впливу*.

1.3 Основні символи змінних та операцій.

Поняття масиву даних

Символи змінних.

Для позначення певної змінної будемо використовувати латинські літери X , Y , Z і т.

Для позначення значень певної змінної будемо використовувати позначки типу X_1, X_2, X_3, X_4, X_5 і т.д. або $X_1, X_2, \dots, X_5$. У загальному окреме значення змінної позначатимемо X_i .

Символи операцій.

Статистика має справу з великими масивами даних, з якими слід проводити багато різних операцій. Тому для скороченого запису створили ряд символів, серед яких найчастіше нам буде зустрічатися СУМА (Σ).

Для запису $X_1 + X_2 + X_3 + X_4 + X_5$ пишуть так:

$$\sum_{i=1}^5 X_i.$$

Масив даних.

Під масивом даних ми будемо розуміти набір даних, отриманий у певному дослідженні. Це можуть бути показники інтелекту та віку досліджуваних, типи темпераменту, соціометричні статуси тощо.

ЛЕКЦІЯ 2

СТАТИСТИЧНЕ СПОСТЕРЕЖЕННЯ. ПОНЯТТЯ ПРО РЕПРЕЗЕНТАТИВНІСТЬ. ГРАФІЧНЕ ПРЕДСТАВЛЕННЯ ДАНИХ

1. Статистичне спостереження: види, способи, принципи та помилки.
2. Генеральна сукупність та вибірка. Репрезентативність вибірки.
3. Статистична значимість.
4. Статистичні таблиці та графіки.

2.1 Статистичне спостереження: види, способи, принципи та помилки

Статистичне спостереження – це процес науково організованого планомірного збору даних.

Можна виділити види статистичного спостереження:

1. За часом реєстрації фактів:
 - поточне (систематичний запис фактів по мірі їх виникнення – запис народження дітей, реєстрація шлюбів і т.д.),
 - періодичне (реєстрація фактів через строго визначені проміжки часу),
 - одночасне (реєстрація фактів по мірі їх необхідності в певний момент часу)
2. За кількістю досліджуваних:
 - суцільне (вивчаються всі досліджувані певної сукупності),
 - несучільне:
 - вибіркове (дає характеристику всієї сукупності на основі обстеження її частини),

- спостереження основного масиву (спостерігаються об'єкти, які займають найбільшу питому вагу в досліджуваній сукупності),
- монографічне (детальне і глибоке вивчення окремих одиниць досліджуваної сукупності).

3. За способами статистичне спостереження ділять на:

- безпосереднє (отримання відомостей шляхом особистого огляду, підрахунку і т.д.),
- документальне (отримання інформації з певних документів),
- опитування (отримання інформації зі слів опитуваного).

До принципів статистичного спостереження відносять:

1. ***Формулювання мети дослідження*** – слід визначити мету дослідження, інакше буде зібрано багато зайвої інформації, і мало – необхідної.

2. ***Визначення об'єкта дослідження*** – слід визначити, яке коло явищ досліджується і в якому аспекті.

3. ***Розробка програми дослідження*** – в якій послідовності і які факти вивчатимуться.

Отримані з допомогою статичного спостереження матеріали можуть бути відповідним чином оброблені і на основі обробки можна зробити певні висновки про досліджуваний процес чи явище. Однак, на висновок можуть вплинути помилки і похибки, які виникають при дослідженні.

1. Помилки реєстрації (виникають внаслідок неправильного запису фактів):

- випадкові (описки, недостатнє знання і т.д.),
- систематичні:
- навмисні (коли опитувані чи реєстратор навмисно повідомляють чи записують неправильні дані),
- ненавмисні (зумовлені систематичними випадковими причинами – несправність приладів, втома досліджуваних після робочого дня тощо).

2. Помилки репрезентативності (характерні лише для несуттєвого спостереження – виникають внаслідок того, що склад відібраної для дослідження вибірки не відтворює усієї сукупності досліджуваних).

Виявити помилки допомагає статистичний контроль (арифметичний та логічний).

2.2 Генеральна сукупність та вибірка. Репрезентативність вибірки

Можливості психологічних досліджень обмежені одним суттєвим фактором – вибіркою. **Генеральна сукупність** – це будь-яка група людей, яку психолог вивчає за вибіркою. Теоретично вважається, що об'єм генеральної сукупності необмежений. Практично ж цей об'єм завжди обмежений та може бути різним в залежності від предмету та задач дослідження. **Вибірка** – це сукупність досліджуваних людей. **Вибірка є частиною генеральної сукупності** (рис. 2.1). **Вибіркою** називається будь-яка підгрупа елементів (досліджуваних), яка була виділена з генеральної сукупності для проведення дослідження. Окремий індивід з вибірки називається **респондентом**. **Об'єм вибірки** позначають буквою *n*.



Рисунок 2.1 Співвідношення понять «генеральна сукупність» та «вибірка»

У статистиці розрізняють:

- *малу* ($n < 30$),
- *середню* ($30 < n < 100$),
- *велику* ($n > 1000$) вибірки.

В залежності від охопту дослідженням аудиторії досліджуваних, воно може бути **повним** (досліджується вся генеральна сукупність) **або вибірковим** (досліджується певна вибірка). Вибірки, в свою чергу, поділяються на залежні та незалежні.

Вибірки називаються незалежними (незв'язними), якщо процедура дослідження й отримані результати вимірювання деякої властивості у досліджуваних однієї вибірки не оказують впливу на особливості протікання цього ж дослідження в респондентів іншої вибірки, і навпаки. **Вибірки називаються залежними** (зв'язними), якщо кожен респондент однієї вибірки так чи інакше пов'язаний з єдиним респондентом іншої вибірки)

Приклади незалежних вибірок: студенти I та II курсу, мешканці різних будинків, вчителі та медпрацівники.

Приклади залежних вибірок: вибірка чоловіків та вибірка їх дружин, вибірка батьків та вибірка їх дітей.

Одна і та сама група досліджуваних, на якій психологічне обстеження проводилося двічі (навіть якщо за різними показниками) вважається залежною або зв'язаною вибіркою.

Вимоги до вибірки.

До вибірки ставиться ряд обов'язкових вимог, визначених передусім цілями та задачами дослідження. Основні з них:

- *Однорідність вибірки*
- *Репрезентативність вибірки* – це вибірка, в якій всі основні ознаки генеральної сукупності представлені приблизно з тією ж частотою та в тій самій пропорції, що й в генеральній сукупності.

Якщо вибірка репрезентативна – то виявлені на ній закономірності можна перенести на генеральну сукупність.

2.3 Математичні показники вибірки

При використанні математичної статистики у психологічних дослідженнях можна виділити дві основні групи методів: *методи описової статистики* і *методи статистичних висновків*.

Методи описової статистики використовують для показників однієї змінної (статистика випадкової вибірки) та виявлення взаємозв'язків між двома та більше змінними (кореляційний аналіз). Описова статистика дає змогу отримати нову інформацію, швидше зрозуміти і всебічно оцінити її.

Методи статистичних висновків надають можливість оцінити похибки, які виникають у процесі статистичних досліджень (статистичне оцінювання); узагальнити параметри генеральної сукупності, отримані на підставі вибірових статистик (перевірка статистичних гіпотез).

Найпоширеніші методи описової статистики (математичні показники вибірки):

- емпіричний розподіл частот;
- міри центральної тенденції (мода, медіана, середнє арифметичне);
- міри мінливості (дисперсія та стандартне відхилення).

2.3.1 Емпіричний розподіл частот

Частоти в математичній статистиці показують як часто (з якою частотою) зустрічаються однакові значення змінної у вибірці.

Наприклад: розглянемо таблицю 2.1 з даними про рівень успішності за 5-бальною шкалою студентської групи з 10 осіб.

Як видно з таблиці, оцінку «5» мають 4 студенти, оцінку «4» – 3 студенти, оцінку «3» – 2 студенти, оцінку «2» – 1 студент; оцінку «1» – 0 студентів. Тобто можна сказати, що «п'ятірка» зустрічається в групі (вибірці) з частотою 4, «четвірка» – з частотою 3 і т.д. Кількість студентів з однаковою

оцінкою складає абсолютну частоту змінної «успішність навчання». Сума всіх абсолютних частот n_i змінної на вибірці повинна дорівнювати об'єму вибірці (загальній кількості респондентів).

Таблиця 2.1

№ студента	Оцінка	№ студента	Оцінка
1	5	6	4
2	5	7	3
3	4	8	5
4	5	9	2
5	3	10	4

Кількість студентів з однаковою оцінкою або абсолютну частоту n_i можна перевести у відсотки наступним чином:

$$\frac{\text{абсолютна частота}}{\text{об'єм вибірки}} \cdot 100\%$$

Тобто, якщо «п'ятірку» отримали 4 студенти, то у відсотках це буде складати:

$$\frac{4}{10} \cdot 100\% = 0,4 \cdot 100\% = 40\%.$$

Абсолютна частота, виражена у відсотках відносно загальної кількості респондентів, називається *відносною частотою*. Відносна частота може також виражатися не у відсотках а десятковим дробом (часткою від одиниці). Наприклад $40\% = 0,4$. Сума всіх відносних частот змінної на вибірці повинна становити 100% (або 1 у випадку вираження десятковим дробом).

Таким чином, на основі аналізу прикладу, можна сформулювати наступні означення.

Емпіричний розподіл частот – математична модель об'єкта реальності у вигляді співвідношення значень змінної x_i і частот їх появи n_i .

Абсолютна частота n_i – кількість об'єктів з однаковим значенням властивості.

Відносна частота – відношення абсолютних частот n_i до загальної кількості об'єктів n .

Для ефективної роботи з частотами змінної на вибірці зручно складати таблиці абсолютних та відносних частот. Для прикладу наведеного вище вона виглядає наступним чином (табл. 2.2):

Таблиця 2.2

Оцінка	Кількість студентів (абсолютна частота)	% від загальної кількості (відносна частота)
5	4	40%
4	3	30%
3	2	20%
2	1	10%
1	0	0%

Для розрахунку абсолютних частот в програмі MS Excel використовуються вбудована функція COUNTIF (знаходяться в блоці статистичних функцій). Розглянемо особливості застосування кожної з них.

Функція COUNTIF підраховує кількість однакових значень в заданому діапазоні клітинок. Наприклад, COUNTIF(B2:B11;5) підраховує кількість клітинок в діапазоні B2:B11, в яких записано число 5. COUNTIF(B2:B11; “низький”) підраховує кількість клітинок в діапазоні B2:B11, в яких записано текст «низький».

Відносна частота розраховується за формулою введеною вручну (клітинка з абсолютною частотою поділити на об’єм вибірки та помножити на 100%. Формат комірки необхідно змінити на відсотковий (права кнопка миші – Формат комірки).

В якості критерію функції COUNTIF в лапках можна також зазначати правило, на основі якого значення комірки враховується або опускається при підрахунку. Наприклад, критерій «=5» означає, що комірка буде рахуватися, якщо її значення дорівнює 5. Замість знаку рівності можна використовувати знак «більше» (>), «менше або дорівнює» (<=), «не дорівнює» (< >). Якщо в якості критерію використовується символ або група символів (слово, скорочення), вони

записуються в лапках в такому вигляді, в якому містяться у комірках. У такому випадку буде підрахована кількість комірок, що містять даний набір символів.

Приклад 2.1. За показником вербальної агресії методики Баса-Дарки отримані наступні тестові бали на вибірці 10 осіб. Необхідно перевести дані в рівні згідно з ключем тесту (в номінативну шкалу), після чого підрахувати кількість респондентів з кожним з трьох рівнів та виразити результати у відсотках (тобто побудувати таблицю абсолютних і відносних частот).

Так, вихідні дані містяться у таблиці 2.3:

Таблиця 2.3

№ студента	Бал тесту	№ студента	Бал тесту
1	14	6	28
2	26	7	27
3	10	8	5
4	15	9	13
5	18	10	7

Ключ до тесту:

0-15 б. – низький рівень;

16-25 б. – середній рівень (норма);

26-... б. – високий рівень.

Для розглянутого прикладу (табл. 2.4) масив інтервалів 16; 25, оскільки ключем виділяється три інтервали:

1. менше 16 (=COUNTIFS(B\$13:B22;><16»));
2. від 16 до 25
(=COUNTIFS(B\$13:B23;>=16»;B\$13:B23;>=25»));
3. більше або дорівнює 26
(=COUNTIFS(B\$13:B22;>=26»)).

Таблиця 2.4

12	бал тесту		інтервал		таблиця частот	
13	1	14	15	низька	6	60%
14	2	26	25	середня	1	10%
15	3	10		висока	3	30%
16	4	15				
17	5	18				
18	6	28				
19	7	27				
20	8	5				
21	9	13				
22	10	7				

2.3.2 Міри центральної тенденції

Міри центральної тенденції (МЦТ) – чисельні показники типових властивостей емпіричних даних. Ці показники дають відповіді на те, наприклад, «який середній рівень інтелекту дітей», «яка в цілому оцінка активності, відповідальності певної групи осіб» тощо. До мір центральної тенденції належать мода, медіана, середнє арифметичне.

Мода (Mo) – значення, яке зустрічається серед емпіричних даних найчастіше. У дискретному ряду **Mo** визначається без обчислення, як значення ознаки з найбільшою частотою.

Під час розрахунку моди може виникнути кілька ситуацій:

1. Два значення ознаки, що стоять поруч, зустрічаються однаково часто. У цьому випадку мода дорівнює середньому арифметичному цих двох значень. Наприклад, у такому ряду даних:

12, 13, 14, 14, 14, 16, 16, 16, 18, 19

$$Mo = (14+16)/2 = 15.$$

2. Два значення, що зустрічаються також однаково часто, але не стоять поруч. У цьому випадку кажуть, що ряд даних має дві моди, тобто він бімодальний.

3. Якщо всі значення даних зустрічаються однаково часто, то кажуть, що ряд не має моди.

Медіана (Md) – значення, яке перебуває на середині упорядкованої послідовності емпіричних даних. Для непарної кількості даних медіану визначають середнім елементом. Якщо кількість значень даних є парною, то медіаною є середнє арифметичне значення центральних сусідніх елементів.

Наприклад, для масиву емпіричних даних 3; 7; 3; 8; 10; 12; 10; 3; 6; 15 мода буде дорівнювати 3, оскільки число 3 зустрічається частіше за інші. Можна також говорити, що мода дорівнює тому значенню, яке має найбільшу абсолютну частоту.

Для обчислення медіани необхідно впорядкувати значення з зростання та знайти середній (центральный) елемент:

3; 3; 3; 6; **8**; 10; 10; 12; 15

Медіана буде дорівнювати 8.

Якщо чисел буде парна кількість, то медіаною буде середнє арифметичне двох центральних значень:

3; 3; 3; 6; **8; 10**; 10; 12; 15; 16.

Медіана буде дорівнювати 9.

Середнє арифметичне обчислюється як відношення суми всіх значень змінної до їх кількості.

$$\bar{X} = \frac{x_1 + x_2 + \dots + x_n}{n} = \frac{1}{n} \sum_{i=1}^n x_i.$$

Середнє арифметичне значення ознаки, обчислене для будь-якої групи, **інтерпретується як значення найбільш типове для цієї групи**. Однак бувають випадки, коли подібна інтерпретація неспроможна (у разі, якщо існує велика різниця між мінімальним і максимальним значеннями ознаки).

Для обчислення значення середнього арифметичного, моди, медіани в MS Excel використовуються відповідно вбудовані функції AVERAGEA, MODE, MEDIAN з категорії «Статистичні». При цьому в дужках необхідно зазначити діапазон комірок, в яких містяться емпіричні дані, або виділити діапазон курсором.

Наприклад: =AVERAGEA(B2:B14), =MODE(B2:B14), =MEDIAN(B2:B14)

2.3.3 Міри мінливості

Використовуючи для опису ряду значень ознаки, тільки міру центральної тенденції (МЦТ), можна сильно помилитися в оцінці характеру досліджуваної сукупності. Це добре видно на наступному прикладі. Припустимо, ми вивчаємо середній вік у двох групах, що складаються кожна з 6-ти осіб. Значення ознаки розподілилися так:

1 група – 10, 10, 10, 50, 50, 50, 50

2 група – 30, 30, 30, 30, 30, 30, 30

Підраховавши середнє значення в кожній із груп, отримаємо $\bar{X}_1 = 30$ і $\bar{X}_2 = 30$. Тобто ми отримали однакові значення, тоді як цілком очевидно, що вибірки взято з різних сукупностей. Помилка сталася через розкид значень віку в цих групах.

Існує кілька способів оцінки ступеня розкиду або розсіювання даних. Основними характеристиками розсіювання є: *розмах (R), дисперсія (D), середньоквадратичне (стандартне) відхилення (σ – сигма), коефіцієнт варіації (V)*.

Найпростіший із параметрів розподілу, розмах – це різниця між максимальним і мінімальним значеннями ознаки:

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

Дисперсія показує розкид значень ознаки відносно свого середнього арифметичного значення, тобто наскільки щільно значення ознаки групуються навколо $\bar{X}$; чим більший розкид, тим сильніше варіюються результати випробовуваних у даній групі, тим більші індивідуальні відмінності між випробовуваними:

Дисперсія визначається за формулою:

$$D = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{X})^2$$

Розрахуємо дисперсію для нашого прикладу з оцінками студентів x_i – оцінка певного студента (табл. 2.5); $\bar{X} = 4$ – середнє арифметичне:

Таблиця 2.5

№ студента	Оцінка	$x_i - \bar{X}$	$(x_i - \bar{X})^2$
1	5	1	1
2	5	1	1
3	4	0	0
4	5	1	1
5	3	-1	1
6	4	0	0
7	3	-1	1
8	5	1	1
9	2	-2	4
10	4	0	0
	40	Сума	10

Отже, $D = \frac{10}{10-1} = 1,111$

Більшу наочність щодо розкиду має середньоквадратичне (стандартне) відхилення, оскільки його розмірність відповідає розмірності вимірюваної величини:

$$\sigma = \sqrt{D}$$

Для обчислення значення дисперсії та стандартного відхилення в програмі MS Excel використовуються відповідно вбудовані функції VARA (або VAR.S), STDEVA (або STDEV.P) з категорії «Статистичні». При цьому в дужках необхідно зазначити діапазон комірок, в яких містяться емпіричні дані.

Коефіцієнт варіації взагалі не має розмірності, що дає змогу порівнювати варіативність випадкових величин, які мають різну природу:

$$V = \frac{\sigma}{\bar{X}} \cdot 100\%.$$

2.3.4 Характеристики асиметрії та ексцесу

Міра асиметрії – *коефіцієнт асиметрії* (As), що розраховується за формулою

$$As = \frac{\sum_{i=1}^n (x_i - \bar{X})^3}{n \sigma^3}$$

Асиметрія характеризує ступінь асиметричності розподілу. Коефіцієнт асиметрії змінюється від мінус до плюс нескінченності ($-\infty < As < +\infty$), для симетричних розподілів $As = 0$.

Міра ексцесу (гостровершинності) – коефіцієнт ексцесу (Ex), що розраховується за формулою:

$$Ex = \frac{\sum_{i=1}^n (x_i - \bar{X})^4}{n \sigma^4} - 3.$$

Коефіцієнт ексцесу також змінюється від мінус до плюс нескінченності ($-\infty < Ex < +\infty$), і $Ex = 0$ для нормального розподілу.

Для обчислення значення асиметрії та коефіцієнту ексцесу в MS Excel використовуються відповідно вбудовані функції SKEW та KURT з категорії «Статистичні». При цьому в дужках необхідно зазначити діапазон комірок, в яких містяться емпіричні дані.

Приклад 2.2. У студентів першого та другого курсу було досліджено рівень депресивного розладу за методикою Бека. Зробити порівняльний аналіз, використовуючи методи описової статистики. Результати тестування наведено в таблиці 2.6.

Таблиця 2.6

1 курс	2 курс
30	24
27	17
23	17
22	15
19	15
19	14
18	14
16	13
15	12
14	12
13	11
12	11
12	8
12	8
10	7
10	7
10	4
10	0

1. Оцінка центральної тенденції:

	1 курс	2 курс
мода	10	17
медіана	14,5	12
середнє	16,22	11,61

Медіана і середнє арифметичне значення на першому курсі вищі, ніж на другому, з чого можна зробити висновок, що рівень депресивного розладу на 1 курсі перевищує рівень розладу на другому курсі. Однак мода на 2-му курсі значно вища, ніж на першому, тобто переважають вищі значення рівня депресивного розладу.

2. Оцінка міри розсіювання:

	1 курс	2 курс
Дисперсія	37,0	30,1
Ст. відхилення	6,1	5,48
Коеф. Варіації	37,5	47,19

Дисперсія і стандартне відхилення на першому курсі вищі, ніж на другому, що свідчить про ширший розкид даних і отже, можна зробити висновок про те, що друга вибірка більш однорідна.

3. Асиметрія та ексцес:

	1 курс	2 курс
асиметрія	0,94	0,01
ексцес	0,06	0,89

Перша вибірка (1 курс) має невелику позитивну лівобічну асиметрію (0,9) та незначний позитивний ексцес (0,06). Друга вибірка (2 курс) також має невелику позитивну асиметрію (0,5) і позитивний ексцес (0,64).

2.4 Графічне зображення частотних розподілів

Таблиця емпіричного розподілу дає майже повну інформацію про вихідну сукупність даних. Проте для наочного сприйняття табличні дані не дуже зручні. Для швидшого сприйняття основних властивостей частотного розподілу використовують графічні методи.

За допомогою графічного зображення розподілу можна як підкреслити найважливіші властивості вихідної сукупності, так і спотворити уявлення про них. При цьому спотворення часто здійснюють цілком свідомо з метою ввести в оману читачів звіту про досліджувані дані.

Наведемо два методи графічного зображення варіаційного ряду: за допомогою полігону, гістограми і кумуляти частот.

2.4.1 Полігон частот

Нехай за варіаційним рядом побудовано його емпіричний розподіл.

<i>Варіанта</i>	x_1	x_2	...	x_m
<i>Частота</i>	n_1	n_2	...	n_m
<i>Відносна частота</i>	w_1	w_2	...	w_m

Полігоном частот називають ламану, відрізки якої послідовно з'єднують між собою точки

$$(x_1, n_1), (x_2, n_2), \dots, (x_m, n_m).$$

Так само полігоном відносних частот називають ламану, відрізки якої послідовно з'єднують між собою точки

$$(x_1, w_1), (x_2, w_2), \dots, (x_m, w_m).$$

Приклад 2.3. Нехай у групі з 27 осіб здійснено опитування, одним з питань якого було таке: “Як часто ви відвідували кінотеатр минулого року?” За відповідями на це запитання було побудовано варіаційний ряд, варіантами якого

були кількість відвідувань кінотеатрів за рук. Емпіричний розподіл частот цього ряду наведено у вигляді таблиці.

Варіанта x_i	0	1	2	3	4	5	6	7
Частота n_i	0	5	1	8	3	2	2	6

Побудуємо за цим розподілом полігон частот. Для цього на вісі абсцис відкладемо значення варіант x_i , на осі ординат – значення відповідних їм частот n_i . Послідовно з'єднавши точки (x_i, n_i) відрізками, отримаємо ламану, що і є полігоном частот (рис 2.2).

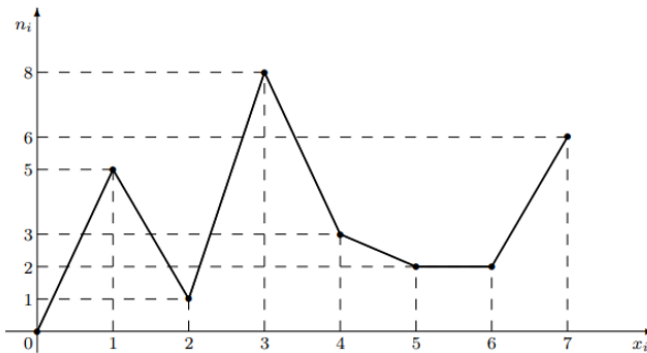


Рисунок 2.2

2.4.2 Гістограма частот

Гістограмою частот емпіричного розподілу досліджуваного варіаційного ряду називають ступінчасту фігуру, що складається з прямокутників. Для кожної варіанти x_i з ненульовою частотою у прямокутній системі координат будують прямокутник висотою n_i , основа якого лежить на осі абсцис так, щоб значення варіанти x_i розміщувались посередині основи.

Гістограма відносних частот будується подібно. При цьому висотами прямокутників є значення відносних частот ω_i .

Зауважимо, що полігон (відносних) частот з'єднує середини вершин прямокутників гістограми (відносних) частот.

Далі на рисунку 2.3 зображено гістограму частот для розподілу з прикладу 2.3 (тонка лінія — полігон частот)

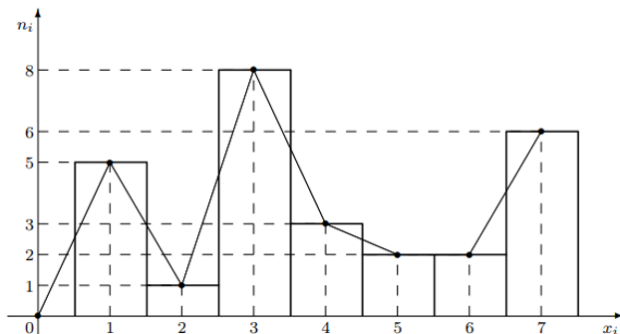


Рисунок 2.3

2.4.3 Вплив графічних зображень на уявлення про сукупність

Графічні зображення емпіричного розподілу вихідної сукупності можуть не лише допомогти краще зрозуміти її основні властивості, а й спотворити уявлення про сукупність.

Приклад 2.4 Опитано 1800 сімей киян про кількість назв газет, які вони періодично купують. У результаті аналізу анкет отримано такий частотний розподіл цього показника.

Варіанта (кількість назв газет)	0	1	2	3
Частота (кількість сімей)	450	500	400	450

Наведемо графічне зображення полігону частот цього розподілу (рис. 2.4 а). Оскільки значення частот перебувають у межах від 400 до 500, полігон розміщується вгорі. Отже, втрачається багато корисної площі!

На вісі ординат передбачають лише значення частот, які зустрічаються в розподілі, тобто локалізують вісь ординат так, щоб графічне зображення повністю використовувало площу

рисунка. При цьому зображення виходить деталізованими завдяки збільшенню масштабу вісі ординат.

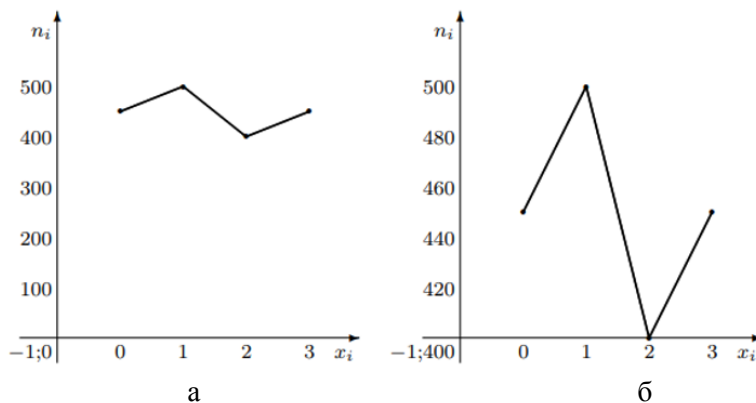


Рисунок 2.4

Для розглядуваного випадку (оскільки $n = 400 \dots 500$) отримуємо інше зображення полігона частот (див. рисунок 2.4 б).

На наш погляд, у розглянутому прикладі частоти різняться неістотно. Справді, вони відхиляються від середнього рівня щонайбільше на 10 %. Тому вважаємо, що в цій ситуації рисунок б швидше спотворює уявлення про розподіл, ніж допомагає зрозуміти його властивості.

ЛЕКЦІЯ 3

ОСНОВНІ СТАТИСТИЧНІ РОЗПОДІЛИ

1. Різновиди форм емпіричних розподілів.
2. Нормально розподіли.
3. Побудова кривої нормального розподілу за емпіричними даними.
4. Перевірка нормальності розподілу результативної ознаки.

3.1 Різновиди форм емпіричних розподілів

Наведемо приклади, що свідчать про розмаїття форм кривих частот розподілів.

Приклад 3.1 Припустимо, що в деякій країні є чотири національних телеканали. Нехай деяка соціологічна служба здійснила опитування 1500 респондентів з тим, щоб з'ясувати, якому національному телеканалу вони віддають перевагу. Дана дослідження наведено в таблиці 3.1.

Таблиця 3.1

Телеканал	1-й	2-й	3-й	4-й
Кількість респондентів, які віддають йому перевагу n_i	371	390	357	382

Побудуємо полігон частот цього розподілу (рис. 3.1)

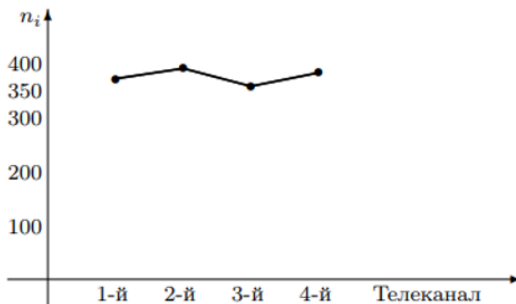


Рисунок 3.1

Як бачимо, полігон частот розглядуваного розподілу мало відрізняється від прямої лінії. Це пов'язано з тим, що уподобання щодо телеканалів опитаної аудиторії розділилися майже порівну.

Приклад 3.2 Розглянемо ідеальну, з позицій статистики, криву частот. Наведемо дані про зріст випадково вибраних 100 десятирічних хлопчиків (табл. 3.2).

Таблиця 3.2

<i>Зріст X, дюйм</i>	52	53	54	55	56	57	58	59	60
<i>Кількість</i>	1	3	10	24	30	22	7	2	1

Середній рівень вихідної сукупності становить 56 дюймів. При цьому крива частотного розподілу симетрична щодо цього рівня. Крім того, хлопчиків, зріст яких значно відрізняється від середнього рівня, небагато і їх кількість швидко зменшується при значному відхиленні від середнього рівня (рис. 3.2). Розподіли схожої форми часто називають куполоподібними.

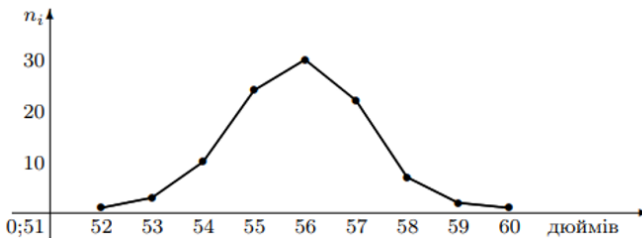


Рисунок 3.2

Приклад 3.3 Наведемо дані дослідження кількості хлопчиків у групі з 1000 сімей, що мають п'ятьох дітей.

<i>Кількість хлопчиків</i>	0	1	2	3	4	5
<i>Кількість сімей</i>	29	140	305	322	167	37

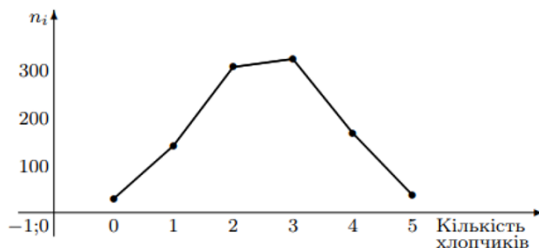


Рисунок 3.3

Як бачимо (рис. 3.3), за формою полігон, як і у прикладі, куполоподібний. Проте наведений розподіл має іншу природу. Вона пов'язана зі структурою досліджуваної ознаки. Маємо властивість дітей “стать”, яка може приймати лише два значення – “хлопчик” або “дівчинка”, і вивчаємо ознаку, яка полягає в розрахунку кількості появ одного з цих значень.

3.2 Нормальні розподіли

Чи не найважливішими в математичній статистиці є так звані нормальні розподіли. З різних позиції нормальна форма вважається ідеальною для емпіричних частотних розподілів. Крім того, нормальні розподіли виникають у низці теоретичних результатів, які є основою багатьох прикладних методів досліджень.

Нормальний розподіл характеризує такі випадкові величини, на які впливає велика кількість різноманітних чинників, причому сила впливу одного окремо взятого чинника значно менша за суму впливів інших чинників. В результаті виходить, що частіше спостерігаються деякі середні значення вимірюваного параметра, рідше крайні, і чим сильніше відрізняється якесь значення від середнього, тим рідше воно зустрічається. Багато біологічних параметрів розподілені подібним чином (зріст, вага тощо).

Психологи вважають, що більшість психологічних властивостей, якостей (інтелект, властивості особистості тощо)

також має нормальний розподіл, саме з цієї посилки виходять під час проведенні стандартизації тестових методик.

Параметри розподілу – це його числові характеристики, що вказують, де «в середньому» розташовуються значення ознаки, наскільки ці значення мінливі та чи спостерігається переважна поява певних значень ознаки. Найбільш практично важливими параметрами є математичне сподівання (μ), дисперсія (D), стандартне відхилення (σ), показники асиметрії та ексцесу.

Основні властивості нормальних розподілів

Форма і властивості нормального розподілу визначаються двома параметрами: його математичним сподіванням μ та стандартним відхиленням σ . Тому нормальний розподіл з параметрами μ і σ позначають $N(\mu, \sigma)$.

Опишемо основні властивості нормального розподілу $N(\mu, \sigma)$.

1. Математичне сподівання розподілу $N(\mu, \sigma)$ дорівнює μ , а його стандартне відхилення σ . Ці два параметри повністю визначають розподіл $N(\mu, \sigma)$.

2. Мода і медіана нормального розподілу дорівнюють μ , тобто збігаються з математичним сподіванням. Фактично можна вважати μ середнім рівнем нормально розподіленої сукупності.

3. Крива нормального розподілу симетрична відносно вертикальної прямої $x = \mu$ і має куполоподібну форму (рис. 3.4)

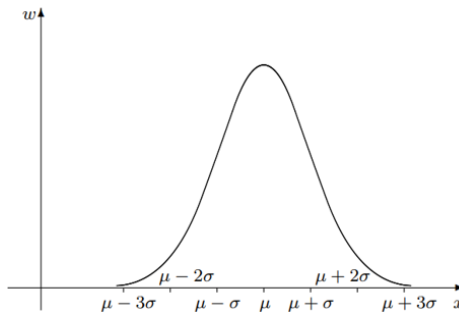


Рисунок 3.4

4. Ймовірність потрапляння випадково взятої варіанти до інтервалу $(\mu - \sigma; \mu + \sigma)$ становить 68,3 %, а до інтервалу $(\mu - 2\sigma; \mu + 2\sigma)$ – 95,4 %.

5. Ймовірність потрапляння випадково взятої варіанти до інтервалу $(\mu - 3\sigma; \mu + 3\sigma)$ становить 99,7 %, тобто на відстані 3σ від μ розміщуються майже всі варіанти нормального розподілу. Цю властивість називають правилом трьох сигм.

6. Середнє абсолютне відхилення нормального розподілу дорівнює 0,8 його стандартного відхилення, тобто $0,8\sigma$. З огляду на це можна формулювати властивості нормального розподілу в термінах лінійних, а не квадратичних відхилень.

7. Крива нормального розподілу задається формулою

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$$

де $f(x)$ – відносні частоти появи кожного конкретного значення випадкової величини x .

У тих випадках, коли які-небудь причини сприяють частішій появі значень, які вищі або, навпаки, нижчі за середнє, утворюються асиметричні розподіли. За лівобічної, або позитивної, асиметрії в розподілі частіше трапляються нижчі значення ознаки, а за правобічної, або від'ємної – вищі значення ознаки. Для симетричних розподілів $As = 0$;

У тих випадках, коли будь-які причини сприяють переважній появі середніх або близьких до середніх значень, утворюється розподіл із позитивним ексцесом. Якщо ж у розподілі переважають крайні значення, причому одночасно і нижчі, і вищі, то такий розподіл характеризується від'ємним ексцесом і в центрі розподілу може утворитися впадина.

У розподілах із нормальною опуклістю $Ex = 0$. Параметри розподілу виявляється можливим визначити тільки по відношенню до даних, представлених принаймні в інтервальній шкалі. Параметри розподілу не враховують істинної психологічної нерівномірності секунд, міліметрів та інших

фізичних одиниць вимірювання. На практиці психолог-дослідник може розраховувати параметри будь-якого розподілу, якщо одиниці, які він використовував під час вимірювання, визнаються розумними в науковому співтоваристві.

3.3 Побудова кривої нормального розподілу за емпіричними даними

Є декілька способів побудови кривої нормального розподілу за емпіричними даними, за умови, що є передумови вважати даний розподіл близький до нормального. За одним з цих способів теоретичні частоти (m') знаходять за формулою

$$m' = \frac{Nh}{\sigma} \cdot \frac{1}{\sqrt{2\pi}} \cdot e^{-\frac{t^2}{2}},$$

де $\frac{1}{\sqrt{2\pi}} \cdot e^{-\frac{t^2}{2}} = \varphi(t)$ – таблична величина, яку шукають за відхиленнями t , а $\frac{Nh}{\sigma}$ – константа, на яку множаться значення $\varphi(t)$ і яка визначає теоретичні частоти виходячи із загальної чисельності одиниць сукупності та числа виділених груп.

План розрахунку теоретичних частот

1. Дані сортуються за зменшенням.
2. Обчислюють частоту повторення кожного значення (абсолютні частоти) це і є емпірична частота n_i .
3. Обчислюють середньо арифметичне вибірки $\bar{X}$.
4. Обчислюємо середньо квадратичне відхилення σ .
5. Знаходимо нормоване відхилення кожного варіанта від середньо арифметичного

$$t = \frac{x - \bar{X}}{\sigma}.$$

6. За знайденим значенням t за таблицею визначають значення $\varphi(t)$.

7. Обчислюємо константу $\frac{nh}{\sigma}$, де n – об'єм вибірки, а

$$h = \frac{x_{max} - x_{min}}{1 + 3,322 \cdot \lg n}.$$

8. Кожне значення $\varphi(t)$ множимо на константу $\frac{nh}{\sigma}$. Результат множення, після округлення і є шуканою частотою m' теоретичної кривої нормального розподілу.

Приклад 3.4. На групі з 30 добровольців-студентів, був проведений дослід з вивчення окорухової координації. Завдання полягало в тому, щоб вражати запропоновані на дисплеї рухомі мішені. Було пред'явлено 10 послідовностей із 25 мішеней. Побудувати криву розподілу величини, що відображає кількість уражених мішеней.

Результат експерименту:

Таблиця 3.3.

Досл-ні	Кіл-сть уражень
1	12
2	21
3	10
4	15
5	15
6	19
7	17
8	14
9	13
10	11
11	20
12	15
13	15
14	14
15	17

$$\bar{X} = 15,2; \quad n = 15;$$

$$h = \frac{21 - 10}{1 + 3,322 \lg 15} = [\lg 15 = 1,18] = 2,23;$$

Таблиця 3.4.

Знач.	n_i	$x_i - \bar{X}$	$(x_i - \bar{X})^2$	t	$\varphi(t)$	$10,55\varphi(t)$	m'
10	1	-5,2	27,04	-1,64	0,104	1,097	1,1
11	1	-4,2	17,64	-1,32	0,167	1,762	1,8
12	1	-3,2	10,24	-1,01	0,24	2,532	2,5
13	1	-2,2	4,84	-0,69	0,314	3,313	3,3
14	2	-1,2	1,44	-0,38	0,371	3,914	3,9
15	4	-0,2	0,04	-0,06	0,398	4,199	4,2
17	2	1,8	3,24	0,57	0,339	3,576	3,6
19	1	3,8	14,44	1,2	0,194	2,047	2,1
20	1	4,8	23,04	1,51	0,128	1,350	1,4
21	1	5,8	33,64	1,83	0,075	0,791	0,8

$$D = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{X})^2 = \frac{1}{14} \cdot 140,4 = 10,03;$$

$$\sigma = \sqrt{D} = 3,17; \quad \frac{nh}{\sigma} = \frac{15 \cdot 2,23}{3,17} = 10,55$$

3.4 Перевірка нормальності розподілу результативної ознаки

Якщо ми застосовуємо параметричні методи (наприклад, формулу для розрахунку коефіцієнта кореляції Браве-Прсона або дисперсійний аналіз) які слід застосовувати тільки тоді, коли відомо або доведено, що розподіл ознаки є нормальним, то в цьому випадку нам необхідно переконатися в нормальності розподілу результативної ознаки. Нормальність розподілу результативної ознаки можна перевірити шляхом розрахунку показників асиметрії та ексцесу і зіставлення їх із критичними значеннями. Розглянемо застосування методу Є. І. Пустильника на прикладі.

Діяти будемо за таким алгоритмом:

1) розрахуємо критичні значення показників асиметрії та ексцесу за формулами Є. І. Пустильника і зіставимо з ними емпіричні значення;

2) якщо емпіричні значення показників виявляться нижчими за критичні, зробимо висновок про те, що розподіл ознаки не відрізняється від нормального.

Розрахунок емпіричних показників асиметрії та ексцесу здійснюватимемо за формулами даними раніше.

Для розрахунків у таблиці необхідне значення середнього арифметичного, яке обчислюється за формулою:

$$\bar{X} = \frac{x_1 + x_2 + \dots + x_n}{n} = \frac{1}{n} \sum_{i=1}^n x_i = \frac{161}{16} = 10,06.$$

Спочатку зробимо розрахунок проміжних значень, який зручно виконувати поетапно, заносючи дані в таблицю (табл. 3.5).

Таблиця 3.5

№	x_i	$x_i - \bar{x}$	$(x_i - \bar{x})^2$	$(x_i - \bar{x})^3$	$(x_i - \bar{x})^4$
1	11	0,94	0,884	0,831	0,781
2	12	2,94	8,644	25,412	74,712
3	13	1,94	3,764	7,301	14,165
4	9	-1,06	1,124	-1,191	1,262
5	10	-0,06	0,004	0,000	0,000
6	11	0,94	0,884	0,831	0,781
7	8	-2,06	4,244	-8,742	18,009
8	10	-0,06	0,004	0,000	0,000
9	15	4,94	24,404	120,554	595,536
10	14	3,94	15,524	61,163	240,982
11	8	-2,06	4,244	-8,742	18,009
12	7	-3,06	9,364	-28,652	87,677
13	10	-0,06	0,004	0,000	0,000
14	10	-0,06	0,004	0,000	0,000
15	5	-5,06	25,604	-129,554	655,544
16	8	-2,06	4,244	-8,742	18,009
сума	161		102,944	30,468	1725,467

$$D = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{X})^2 = \frac{102,944}{15} = 6,8629;$$

$$\sigma = \sqrt{D} = 2,62.$$

Підставляючи у формули для розрахунку As і Ex отримані значення n , σ і відповідні значення з таблиці, отримуємо:

$$As = \frac{\sum_{i=1}^n (x_i - \bar{X})^3}{n \sigma^3} = \frac{30,468}{16 \cdot 2,62^3} = +0,106;$$

$$Ex = \frac{\sum_{i=1}^n (x_i - \bar{X})^4}{n \sigma^4} - 3 = \frac{1725,467}{16 \cdot 2,62^4} - 3 = -0,711.$$

Тепер розрахуємо критичні значення для показників As і Ex за формулами Є. І. Пустильника:

$$As_{кр} = 3 \cdot \sqrt{\frac{6 \cdot (n-1)}{(n+1)(n+3)}};$$

$$Ex_{кр} = 5 \cdot \sqrt{\frac{24n(n-2)(n-3)}{(n+1)^2(n+3)(n+5)}},$$

де n кількість випробувань.

У даному випадку маємо:

$$As_{кр} = 3 \cdot \sqrt{\frac{6 \cdot (16-1)}{(16+1)(16+3)}} = 3 \cdot \sqrt{\frac{6 \cdot 15}{17 \cdot 19}} =$$

$$= 3 \cdot \sqrt{\frac{90}{323}} = 1,58;$$

$$Ex_{кр} = 5 \cdot \sqrt{\frac{24 \cdot 16(16-2)(16-3)}{(16+1)^2(16+3)(16+5)}} =$$

$$= 5 \cdot \sqrt{\frac{24 \cdot 16 \cdot 14 \cdot 13}{17^2 \cdot 19 \cdot 21}} = 3,89.$$

$$As_{\text{емп.}} = 0,106; \quad As_{\text{емп.}} < As_{\text{кр.}}$$

$$Ex_{\text{емп.}} = -0,711; \quad Ex_{\text{емп.}} < Ex_{\text{кр.}}$$

Оскільки емпіричні значення As і Ex менші за критичні значення, то можна зробити такий висновок: розподіл результативної ознаки в цьому прикладі не відрізняється від нормального розподілу.

ЛЕКЦІЯ 4

СТАНДАРТНИЙ НОРМАЛЬНИЙ РОЗПОДІЛ. ДЕЯКІ СПЕЦІАЛЬНІ РОЗПОДІЛИ

1. Стандартний нормальний розподіл.
2. Спеціальні розподіли.

4.1 Стандартний нормальний розподіл

Раніше вже зауважувалося, що все нормально розподіли мають подібне властивості. Фактично нормальний розподіл визначається двома параметрами: математичним сподіванням μ і стандартним відхиленням σ . Тому вивчення всього класу нормальних розподілів природно звести до одного стандартного випадку.

Середнє арифметичне значень дискретної випадкової величини, які спостерігаються в дослідах, прямує до її математичного сподівання

Зокрема, *найважливіша практична задача*, що пов'язана з вивченням розподілів, полягає у визначення ймовірності потрапляння досліджуваного параметра до деякого інтервалу (a, b) можливих значень. У контексті сформульованої ідеї це означає, що хотілося б знати, як розв'язання цієї задачі загалом можна звести до її розв'язання для деякого стандартного розподілу з класу нормальних.

4.1.1 Стандартний нормальний розподіл

Таким стандартним розподілом у класи нормальних вважають розподіл $N(0,1)$ з математичним сподіванням 0 та стандартним відхиленням 1.

Означення. Розподіл $N(0,1)$ називають *стандартним нормальним розподілом*.

Параметр, розподілений за стандартним нормальним законом, прийнято позначати z , а параметр, розподілений за довільним нормальним законом, — x .

Крива стандартного нормального розподілу задається формулою

$$f(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}}$$

Графічно криву стандартного нормального розподілу представлено на рисунку 4.1

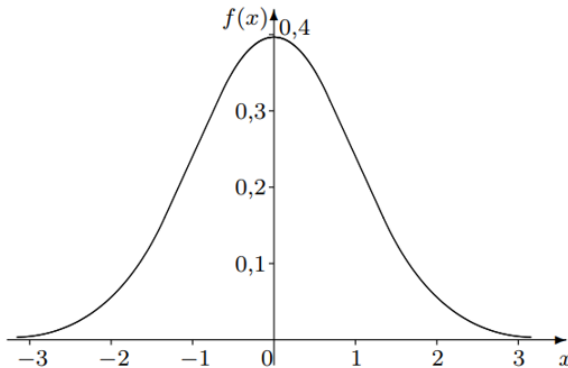


Рисунок 4.1

Виявляється, є проста формула зв'язку між кривою стандартного нормального розподілу та кривою довільного розподілу $N(\mu, \sigma)$.

Теорема Для того щоб перейти від кривої частот нормально розподіленого за законом $N(\bar{X}, \sigma)$ параметра x до кривої частот нормально розподіленого за законом $N(0,1)$ параметра z , необхідно здійснити таку заміну змінних:

$$z = \frac{x - \mu}{\sigma} \quad (*)$$

За формулою $(*)$ для обчислення ймовірності потрапляння параметра x до інтервалу (a, b) достатньо обчислити ймовірність потрапляння параметра z до інтервал

$$\left(\frac{a - \mu}{\sigma}, \frac{b - \mu}{\sigma}\right).$$

Отже, необхідно навчитися визначати ймовірність потрапляння стандартного параметра z до заданого інтервалу.

Означення. Функція Лапласа $\Phi(t)$ є означеною для $t > 0$ і дорівнює ймовірності потрапляння розподіленого за стандартним нормальним законом параметра z до інтервалу $(0, t)$.

Іншими словами, значення функції Лапласа $\Phi(t)$ дорівнює площі фігури, обмеженої зверху кривою стандартного нормального розподілу, знизу – віссю абсцис, а по краях – відрізками прямих $z = 0$ та $z = t$. Для зручності використання значення функції Лапласа протабульовані. Як використовувати функцію Лапласа покажемо на прикладі.

Приклад 4.1. Нехай параметр x нормально розподілений за законом $N(8,2; 2,5)$, тобто математичне сподівання μ його розподілу дорівнює 8,2, а стандартне відхилення $\sigma = 2,5$. Обчислимо ймовірність потрапляння параметра до інтервалу:

1) $(4; 15)$

2) $(10; 12)$

Розглянемо випадки окремо. Насамперед необхідно перейти до стандартного параметра z , використовуючи формулу $z = \frac{x - \mu}{\sigma}$:

$$\frac{4 - 8,2}{2,5} = -1,68; \quad \frac{15 - 8,2}{2,5} = 2,72.$$

Отже, необхідно визначити ймовірність потрапляння стандартного параметра z до інтервалу $(-1,68; 2,72)$. Але, як зауважувалося, ця ймовірність дорівнює площі фігури, обмеженої зверху кривою стандартного нормального розподілу, знизу віссю абсцис, а по краях відрізками прямих $z = -1,68$ та $z = 2,72$. Цю фігуру можна розділити на дві частини: перша від $z = -1,68$ до $z = 0$, друга від $z = 0$ до $z = 2,72$. Тоді шукана площа дорівнює сумі площ цих двох частин.

Отже, шукана ймовірність дорівнює сума ймовірностей потрапляння стандартного параметра z до інтервалів $(-1,68; 0)$ та $(0; 2,72)$. За означенням функції Лапласа друга з цих ймовірностей дорівнює $\Phi(2,72)$. Для того щоб визначити ймовірність потрапляння до першого інтервалу, зауважимо, що крива стандартного нормального розподілу симетрична відносно осі ординат $z = 0$. Тому ймовірність потрапляння параметра z до інтервалу $(-1,68; 0)$ дорівнює ймовірності його потрапляння до $(0; 1,68)$, а отже, значенню функції Лапласа $\Phi(1,68)$.

Остаточно, скориставшись даними таблиці для функції Лапласа, отримаємо, що шукана ймовірність дорівнює

$$\Phi(1,68) + \Phi(2,72) = 0,4535 + 0,4967 = 0,9502 = 95,02 \%$$

2. За формулою перейдемо до стандартного параметра z :

$$\frac{10 - 8,2}{2,5} = 0,72; \quad \frac{12 - 8,2}{2,5} = 1,52.$$

Необхідно визначити ймовірність потрапляння стандартного параметра z до інтервалу $(0,72; 1,52)$, тобто площу фігури, обмеженої зверху кривою стандартного нормального розподілу, знизу віссю абсцис, а по краях відрізками прямих $z = 0,72$ та $z = 1,52$. Цю площу можна знайти, віднявши від площі фігури між $z = 0$ і $z = 1,52$ площу фігури між $z = 0$ і $z = 0,72$.

Отже, шукана ймовірність за даними таблиці для функції Лапласа дорівнює

$$\Phi(1,52) - \Phi(0,72) = 0,4357 - 0,2642 = 0,1715 = 17,15 \%.$$

4.1.2 Біноміальні розподіли

Припустимо, вивчається ознака, яка може набувати два можливих значення. Нагадаємо, що шкали вимірювання таких ознак було названо дихотомічними. Нехай одне з цих двох значень “важливіше” від іншого. Припустимо також, що з вихідної сукупності даних за певним критерієм вибрано деякі підгрупи однакового обсягу. Тоді можна розглянути похідну від початкової ознаку. Зокрема, у кожній підгрупі можна розрахувати кількість “важливих” значень. Іншими словами,

підгрупи утворюють нову сукупність, яка вивчається за ознакою “кількість “важливих” значень у підгрупі”.

Пояснимо наведені міркування на прикладах. У прикладі, який ми розглядали раніше, аналізувалась ознака “кількість хлопчиків у сім’ях з п’ятьма дітьми”. Головною ознакою для неї була “стать дитини”, “важливим” значенням “хлопчик”, вихідною сукупністю – сукупність усіх дітей, вибраними підгрупами в ній – діти сімей з п’ятьма дітьми.

Наведемо інший приклад. Розглянемо дихотомічну ознаку, утворену двома можливими значеннями “потяг прийшов на кінцеву станцію із запізненням” та “потяг прийшов на кінцеву станцію без запізнення”. Інтерес може становити, наприклад, ознака “кількість запізнень за день”, і цю ознаку можна вивчати тривалий час.

Виявляється, емпіричні розподіли подібних сукупностей часто добре описуються такою моделлю. Нехай маємо групи незалежних даних обсягом n . Припустимо, що ці дані можуть приймати лише два можливих значення – “важливе” та “неважливе” і “важливе” значення зустрічається з теоретичною відносною частотою p . Розподіл кількості “важливих” значень у групах називають *біноміальним* з параметрами $(p; n)$.

Наведемо основні властивості біноміального розподілу з параметрами $(p; n)$.

1. Варіантами біноміального розподілу є можлива кількість “важливих” значень у групі з n даних, тобто числа $0, 1, 2, \dots, n$. Тому біноміальний розподіл дискретний.

2. Середнє значення μ біноміального розподілу дорівнює np .

3. Стандартне відхилення біноміального розподілу

$$\sigma = \sqrt{np(1 - p)}.$$

4. Якщо значення p мале, то біноміальний розподіл має лівосторонню асиметрію (рис. 4.2 а). Якщо значення $p \approx 50\%$, то полігон біноміального розподілу має куполоподібну форму

(рис. 4.2 в). Якщо значення p близьке до 100 %, то біноміальний розподіл має виразну правосторонню асиметрію (рис. 4.2 б).

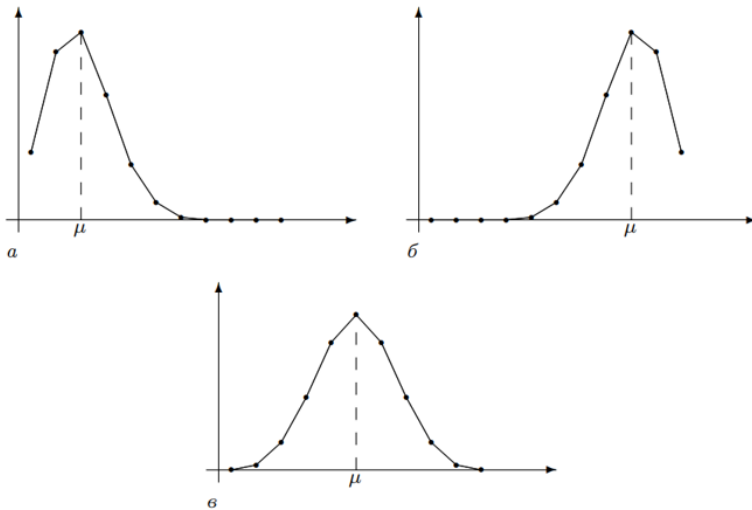


Рисунок 4.2

4.1.3 Пуассонові розподіли

Раніше було розглянуто приклади розподілів із сильною лівосторонньою асиметричністю. У цих розподілах переважають невеликі значення ознаки і рідко зустрічаються великі. Такі розподіли часто вдається описати за допомогою теоретичного розподілу Пуассона. Опишемо основні властивості розподілу Пуассона.

1. Розподіл Пуассона повністю визначається єдиним параметром, який позначають λ .
2. Розподіл Пуассона дискретний, його варіанти – це ціла невід’ємні числа $0, 1, 2, 3, \dots$.
3. Математичне сподівання та дисперсія розподілу Пуассона збігаються і дорівнюють його параметру λ .
4. Відносна частота варіанти m розподілу Пуассона

$$w(m) = \frac{\lambda^m e^{-\lambda}}{1 \cdot 2 \cdot 3 \cdot \dots \cdot m}$$

Полігон відносних частот розподілу Пуассона має дві типові форми: при малих та великих значеннях λ . Їх приклади наведено на рисунку 4.3 а розподіл Пуассона з $\lambda = 2$, б) з $\lambda = 7$.

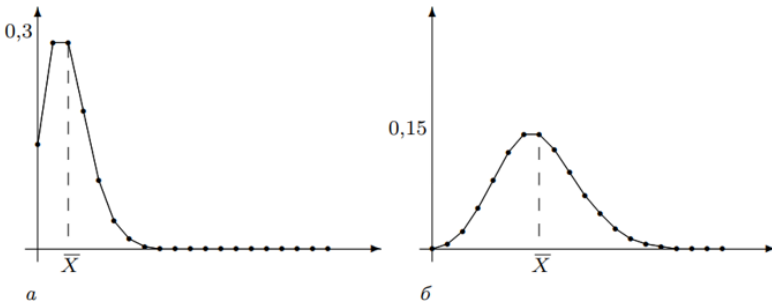


Рисунок 4.3

4.1.4 U-подібні розподіли

Розглянемо розподіл статистичних даних U-подібної форми. Найпростішим теоретичним розподілом, крива якого має U-подібну форму (рис. 4.4), є так званий арксинус-розподіл. Крива цього розподілу задається формулою

$$f(x) = \frac{1}{\pi \sqrt{x(1-x)}}$$

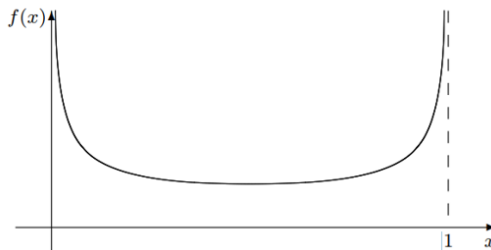


Рисунок 4.4

На жаль, допоки не існує простих стандартних статистичних розподілів, які б точніше описували таку форму і були достатньо загальними. Як правило, у кожному конкретному випадку потрібно підбирати процедуру, за допомогою якої можна згенерувати модель розподілу для даних з полігоном U-подібної форми.

Найчастіше такі процедури полягають у такому: припускається, що деяка частина сукупності розподілена за законом з лівосторонньою асиметрією, а інша з правосторонньою. Далі ці два розподіли частин сукупності об'єднують в один. Це роблять за допомогою так званої операції згортки. У контексті зробленого припущення ця операція є зваженою пропорціями частин сукупності сумою двох розподілів.

Висновки

У процесі дослідження емпіричної сукупності даних намагаються не лише обчислити їх найважливіші числові характеристики, а й описати ці дані за допомогою деякої теоретичної моделі. Для цього найчастіше використовують такий підхід.

1. Спочатку будують емпіричний розподіл відносних частот досліджуваної сукупності, а потім за його зовнішнім виглядом або за природою даних підбирають один з класичних добре вивчених теоретичних розподілів.

2. Перевіряють, чи узгоджуються емпіричний та теоретичний розподіли. Якщо теоретична модель добре описує вихідні емпіричні дані, то кажуть, що теоретичний та емпіричний розподіли поведуться однаково. В такому випадку, вважають, що досліджувана сукупність даних має такі ж властивості, що й теоретичний розподіл.

Вправи

1. Визначити ймовірність потрапляння випадкового параметра, розподіленого за законом $N(7; 2)$, до інтервалу $(0; 5)$.

2. Розглянути біноміальний розподіл з параметрами $p = 0,6$ та $n = 6$. Визначити його математичне сподівання, стандартне відхилення та відносну частоту варіанти 4.

3. Зобразити криву пуассонового розподілу з параметром $\lambda = 4$. Обчислити його математичне сподівання та стандартне відхилення.

4.2 Деякі спеціальні розподіли

Розглянемо три таких розподіли: Стюдента, Пірсона та Фішера-Снедекора.

4.2.1 t-розподіл Стюдента

Розподіл Стюдента найчастіше використовують при статистичному оцінювання та перевірці гіпотез за малих вибірок.

Нехай $X_1, \dots, X_n$ — це незалежні випадкові величини з розподілу $N(\mu, \sigma^2)$, тобто це вибірка розміру n з популяції з нормальним розподілом з математичним сподіванням μ і дисперсією S .

Нехай середнє вибірки та її дисперсія:

$$\bar{X} = \frac{1}{n} \sum_{i=1}^n x_i$$
$$S = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{X})^2$$

Тоді випадкова величина

$$\frac{\bar{X} - \mu}{\sigma / \sqrt{n}}$$

має стандартний нормальний розподіл (тобто, з математичним сподіванням 0 і дисперсією 1), а випадкова величина

$$\frac{\bar{X} - \mu}{S/\sqrt{n}}$$

(де ми підставили S замість σ) має t -розподіл Стюдента з $n - 1$ ступенями вільності. Через те що ми замінили σ на S єдина неспостережувана величина тут це μ отже ми можемо використати це, щоб знайти довірчі інтервали для μ . Зауважте, що незважаючи на те, що вони базуються на тій самій вибірці $X_1, \dots, X_n$ чисельник і знаменник у попередньому виразі – незалежні випадкові величини.

Означення. Припустимо X – сукупність, розподілена за нормальним законом $N(\mu; \sigma)$. Розглянемо довільну просту випадкову вибірку обсягом n . Нехай σ – її виправлене стандартне відхилення. Тоді розподіл

$$t = \frac{\bar{X} - \mu}{S/\sqrt{n}},$$

який формується за всіма простими випадковими вибірками обсягом n , називають *розподілом Стюдента* з $(n - 1)$ ступенями вільності і позначають $t(n - 1)$.

Наведемо властивості розподілу Стюдента.

1. Він симетричний відносно нуля.
2. При великих n поводить як нормальний.
3. Математичне сподівання розподілу з $(n - 1)$ ступенями вільності дорівнює нулю, а дисперсія $n/(n - 2)$.

Криві t -розподілу для деяких n зображено на рисунку 4.5.

Як зазначалося, чи не найважливішим при робота з розподілами є обчислення ймовірностей потрапляння випадково взятого значення сукупності до заданого інтервалу. Насправді при оцінюванні параметрів та перевірці гіпотез необхідно вміти розв'язувати обернену задачу. А саме необхідно вміти визначати такий інтервал, що ймовірність потрапляння випадково взятого значення сукупності до нього дорівнює заданій. При цьому, як правило, інтервал має одну з таких форм: $(-a; a)$ або $(\infty; a)$, а

задана ймовірність приймає одне з кількох стандартних значень. Покажемо на прикладі як розв'язувати цю задачу в разі розподілів Стюдента.

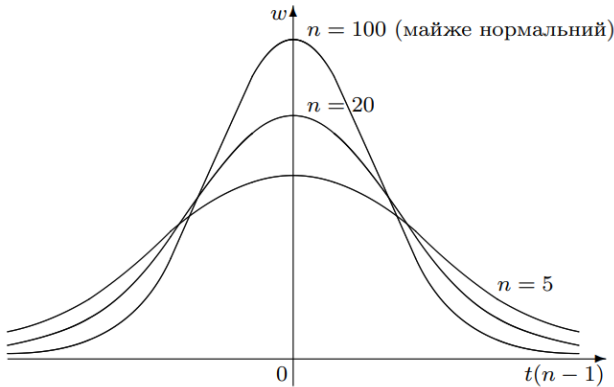


Рисунок 4.5

Приклад 4.2. Припустимо, необхідно визначити інтервал вигляду $(-a; a)$, ймовірність потрапляння до якого розподіленого за законом $t(9)$ випадкового значення дорівнює 95 %. Для цього скористаємося таблицею критичних точок розподілу Стюдента. Заданій ймовірності 95 % відповідає рівень значущості

$$\alpha = 1 - 0,95 = 0,05$$

і двобічна критична область. Отже, розподілене за законом $t(9)$ випадкове значення потрапляє до інтервалу $(-2,26; 2,26)$ з ймовірністю 95 %. Нехай тепер необхідно визначити інтервал вигляду $(-\infty; a)$, ймовірність потрапляння до якого розподіленого за законом $t(150)$ випадкового значення дорівнює 99 %. У цьому разі необхідно взяти рівень значущості $\alpha = 1 - 0,99 = 0,01$ і одnobічну критичну область. Оскільки $150 > 120$, використаємо останній рядок таблиці критичних значень. Отже, розподілене за законом $t(150)$ випадкове значення потрапляє до інтервалу $(-\infty; 2,33)$ з ймовірністю 99 %.

4.2.2 χ^2 – розподіл Пірсона

При визначенні інтервальних оцінок для середнього значення та пропорції у статистиці використовують, як правило, нормальний розподіл і розподіл Стюдента. Розподіл же Пірсона застосовують при інтервальному оцінюванні дисперсії. Формально розподіл Пірсона означають так.

Означення. Позначимо Z генеральну сукупність, розподілену за стандартним нормальним законом $N(0; 1)$. Розглянемо довільну просту випадкову вибірку обсягом n із сукупності Z . Тоді розподіл

$$S_1^2 + S_2^2 + \dots + S_n^2$$

суми квадратів значень вибірки, який формується за всіма простими випадковими вибірками з Z обсягом n , називають розподілом Пірсона з n ступенями вільності. Цей розподіл позначають $\chi^2(n)$.

Наведемо властивості розподілу Пірсона.

1. Його значення перебуває в інтервалі $[0; \infty)$.
2. Для малих значень n розподіл має лівобічну асиметрію. Зі збільшенням n він стає дедалі симетричним і поводить як нормальний. При цьому для $n > 100$ маємо співвідношення

$$\sqrt{2\chi^2(n)} - \sqrt{2n - 1} \approx N(0; 1)$$

3. Математичне сподівання розподілу з n ступенями вільності дорівнює n , а дисперсія становить $2n$.

Криві розподілу Пірсона для різних значень n зображено на рисунку 4.6.

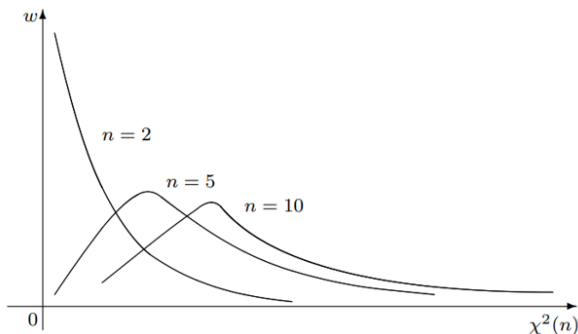


Рисунок 4.6

Розглянемо, як визначити інтервали спеціального вигляду, ймовірність потрапляння до яких дорівнює заданій, для випадку розподілів Пірсона.

Приклад 4.3. Нехай необхідно знайти інтервал вигляду $(a; \infty)$, ймовірність потрапляння до якого розподіленого за законом $\chi^2(15)$ випадкового значення дорівнює 99 %. Для цього скористаємося таблицею критичних точок розподілу Пірсона, де вказано відповідні розглядуваній задачі значення a . Отримаємо, що розподілене за законом $\chi^2(15)$ випадкове значення потрапляє до інтервалу $(5,23; \infty)$ з ймовірністю 99 %. Нехай тепер необхідно знайти інтервал вигляду $(0; a)$, ймовірність потрапляння до якого розподіленого за законом $\chi^2(12)$ випадкового значення дорівнює 95 %. Тоді ймовірність потрапляння до інтервалу $(a; \infty)$ розподіленого за законом $\chi^2(12)$ випадкового значення дорівнює 5 %. Використовуючи таблицю критичних точок розподілу Пірсона, отримуємо значення $a = 21,03$.

Насамкінець розглянемо такий приклад. Нехай необхідно знайти інтервал вигляду $(0; a)$, ймовірність потрапляння до якого розподіленого за законом $\chi^2(220)$ випадкового значення дорівнює 90 %. На жаль, даних про критичні точки розподілу

Пірсона для $n = 220$ ступенів вільності в таблиці немає. Проте можна скористатися формулою

$$\sqrt{2\chi^2(n)} - \sqrt{2n - 1} \approx N(0; 1)$$

Насамперед обчислимо

$$\sqrt{2 \cdot 220 - 1} = \sqrt{439} \approx 20,95.$$

Далі визначимо таке число b , що ймовірність потрапляння до інтервалу $(-20,95; b)$ випадкового значення, розподіленого за стандартним нормальним законом, дорівнює 90 %. Для цього за таблицею значень функції Лапласа спочатку обчислимо

$$\Phi(20,95) \approx 0,5.$$

Для визначення b скористаємося рівністю

$$\Phi(b) = 0,90 - 0,5 = 0,4.$$

За таблицею значень функції Лапласа маємо $b = 1,28$.

Тепер за формулою $\sqrt{2\chi^2(n)} - \sqrt{2n - 1} \approx N(0; 1)$ бачимо, що випадкова величина $2\chi^2(220)$ потрапляє до інтервалу $(0; 1,28 + 20,95) = (0; 22,23)$ з ймовірністю 90 %. Остаточню шукане значення

$$a = \frac{22,23^2}{2} \approx \frac{494,17}{2} \approx 247,08.$$

Отже, розподілене за законом $\chi^2(220)$ випадкове значення потрапляє до інтервалу $(0; 247,08)$ з ймовірністю 90 %.

4.2.3 F-розподіл Фішера-Снедекора

Розглянемо ще один статистичний розподіл, який використовують при порівнянні стандартних відхилень двох незалежних генеральних сукупностей.

Означення. Нехай X, Y — дві незалежні нормально розподілені генеральні сукупності. Припустимо, їх стандартні відхилення дорівнюють σ_X, σ_Y . Розглянемо всеможливі прості

випадкові вибірки обсягом m з X та обсягом n з Y . Нехай s_X і s_Y – їх виправлені стандартні відхилення. Тоді величина

$$\frac{s_X^2 / \sigma_X^2}{s_Y^2 / \sigma_Y^2}$$

має F-розподіл Фішера-Снедекора з $m - 1$ ступенями вільності в чисельнику та з $n - 1$ ступенями вільності у знаменнику. Цей розподіл позначатимемо $F_{m-1, n-1}$.

Наведемо властивості розподілу Фішера-Снедекора.

1. Значення розподілу перебуває в межах від нуля до $+\infty$.
2. Подібно до розподілу Пірсона F-розподіл має правобічну асиметрію.
3. Розподіл наближається до нормального зі збільшенням кількостей ступенів вільності в чисельнику та знаменнику. Криві низки розподілів Фішера-Снедекора зображено на рисунку 4.7.

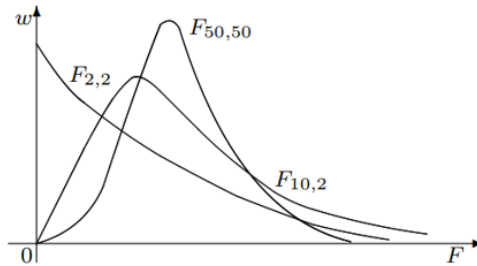


Рисунок 4.7

ЛЕКЦІЯ 5

МІРИ ЗВ'ЯЗКУ МІЖ ОЗНАКАМИ

1. Загальні положення.
2. Коефіцієнт рангової кореляції r_s Спірмена.
3. Коефіцієнт кореляції Браве-Пірсона.
4. Інтерпретація коефіцієнтів кореляції.

5.1 Загальні положення

Аналіз зв'язків між ознаками – головний вид завдань, що зустрічається практично в будь-якому емпіричному дослідженні.

Під час вивчення *кореляцій* намагаються встановити, чи існує якийсь зв'язок між двома показниками в одній вибірці (наприклад, між зростом і вагою дітей або між рівнем IQ і шкільною успішністю) або між двома різними вибірками (наприклад, при порівнянні пар близнюків), і якщо цей зв'язок існує, то чи супроводжується збільшення одного показника зростанням (позитивна кореляція) або зменшенням (негативна кореляція) іншого.

Іншими словами, кореляційний аналіз допомагає встановити, чи можна передбачати можливі значення одного показника, знаючи величину іншого.

Кореляційний зв'язок – це узгоджені зміни двох ознак або більшої кількості ознак (множинний кореляційний зв'язок). Кореляційний зв'язок відображає той факт, що мінливість однієї ознаки перебуває в деякій відповідності з мінливістю іншої.

Кореляційна залежність – це зміни, які вносять значення однієї ознаки у ймовірність появи різних значень іншої ознаки. Стохастична означає ймовірнісна. Зв'язки між випадковими явищами називають ймовірнісними або стохастичними зв'язками. У кореляційних зв'язках кожному значенню однієї ознаки може відповідати певний розподіл значень іншої ознаки, але не певне її значення.

Обидва терміни – кореляційний зв'язок і кореляційна залежність – часто використовують як синоніми. Тим часом, узгоджені зміни ознак і кореляційний зв'язок між ними, що відображає це кореляційний зв'язок між ними може свідчити не про залежність цих ознак між собою, а залежність обох цих ознак від якоїсь третьої ознаки або поєднання ознак, не розглянутих у дослідженні.

Залежність передбачає вплив, зв'язок – будь-які узгоджені зміни, які можуть пояснюватися сотнями причин. Кореляційні зв'язки не можуть розглядатися як свідчення причинно-наслідкового зв'язку, вони свідчать лише про те, що змінам однієї ознаки, як правило, супроводжують певні зміни іншої, але чи знаходиться причина змін в одній з ознак чи вона опиняється за межами досліджуваної пари ознак, нам невідомо.

Кореляційні зв'язки різняться за **формою, напрямом і ступенем (силою)**.

За **формою** кореляційний зв'язок може бути прямолінійним або криволінійним. Прямолінійним може бути, наприклад, зв'язок між кількістю тренувань на тренажері та кількістю правильно розв'язуваних задач у контрольній сесії. Криволінійним може бути, наприклад, зв'язок між рівнем мотивації та ефективністю виконання завдання (рис. 5.1).

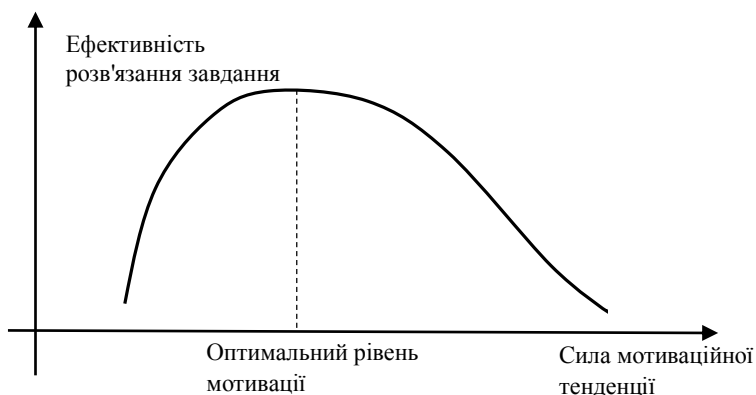
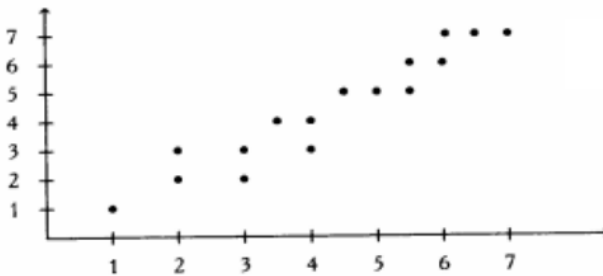


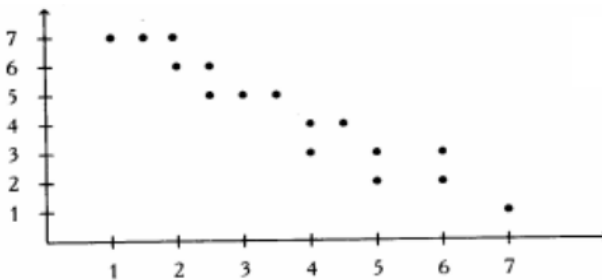
Рисунок 5.1

У разі підвищення мотивації ефективність виконання завдання спочатку зростає, потім досягається оптимальний рівень мотивації, якому відповідає максимальна ефективність виконання завдання; подальше підвищення мотивації супроводжує вже зниження ефективності.

За **напрямком** кореляційний зв'язок може бути **позитивним** («прямим») і **негативним** («зворотним»). За позитивної прямолінійної кореляції більш високим значенням однієї ознаки відповідають більш високі значення іншої, а більш низьким значенням однієї ознаки – низькі значення іншої (рис. 4.2). За негативної кореляції співвідношення зворотні.



а



б

Рисунок 5.2. Схема прямолінійних кореляційних зв'язків:
а – позитивний (прямий) зв'язок, б – негативний (зворотний)
зв'язок

За позитивної кореляції коефіцієнт кореляції має позитивний знак, наприклад $r = +0,207$, за від'ємної кореляції – від'ємний знак, наприклад $r = -0,207$.

Ступінь, сила або тіснота кореляційного зв'язку визначається за величиною коефіцієнта кореляції.

Сила зв'язку не залежить від його спрямованості і визначається за абсолютним значенням коефіцієнта кореляції.

Коефіцієнт кореляції – це величина, яка може варіювати в межах від +1 до -1. У разі повної позитивної кореляції цей коефіцієнт дорівнює плюс 1, а в разі повної негативної – мінус 1. На графіку цьому відповідає пряма лінія, що проходить через точки перетину значень кожної пари даних:

У разі ж якщо ці точки не шикуються по прямій лінії, а утворюють «хмару», коефіцієнт кореляції за абсолютною величиною стає меншим за одиницю і в міру округлення цієї хмари наближається до нуля.

У разі якщо коефіцієнт кореляції дорівнює 0, обидві змінні повністю незалежні одна від одної.

Використовують дві системи класифікації кореляційних зв'язків за їхньою силою: загальна та приватна.

Загальна класифікація кореляційних зв'язків:

- 1) сильний, або тісний при коефіцієнті кореляції $r > 0,70$;
- 2) середня при $0,50 < r < 0,69$;
- 3) помірний при $0,30 < r < 0,49$;
- 4) слабка при $0,20 < r < 0,29$;
- 5) дуже слабка при $r < 0,19$.

Часткова класифікація кореляційних зв'язків:

1) висока значуща кореляція при r , що відповідає рівню статистичної значущості $\rho \leq 0,01$

2) значуща кореляція при r , що відповідає рівню статистичної значущості $\rho \leq 0,05$;

3) тенденція достовірного зв'язку при r , що відповідає рівню статистичної значущості $\rho \leq 0,10$;

4) незначуща кореляція при r , що не досягає рівня статистичної значущості.

Зазвичай прийнято орієнтуватися на другу класифікацію, оскільки вона враховує обсяг вибірки. Водночас необхідно пам'ятати, що сильна, або висока, кореляція – це кореляція з коефіцієнтом $r > 0,70$, а не просто кореляція високого рівня значущості.

В якості міри кореляції використовують:

1) емпіричні міри тісноти зв'язку, багато з яких було отримано ще до відкриття методу кореляції, а саме:

- коефіцієнт асоціації, або тетрахоричний показник зв'язку;
- коефіцієнти взаємної спряженості Пірсона і Чупрова;
- коефіцієнт Фехнера;
- коефіцієнт кореляції рангів;

2) лінійний коефіцієнт кореляції r ;

3) кореляційне відношення η ;

4) множинні коефіцієнти кореляції та ін.

У психологічних дослідженнях найчастіше застосовують коефіцієнт лінійної кореляції r -Пірсона та методи рангової кореляції Спірмена і Кендала. Однак метод Пірсона є параметричним і тому не позбавлений недоліків, властивих параметричним методам (необхідно, щоб дані були виміряні в інтервальних шкалах або розподіл не відрізнявся від нормального). Параметричними є також методи визначення кореляційного відношення та підрахунку множинних коефіцієнтів кореляції.

Метод рангової кореляції Спірмена, є непараметричним методом, він є універсальним і працює з даними виміряними в будь-яких шкалах і простий у застосуванні.

Унікальність методу рангової кореляції полягає в тому, що він дає змогу зіставляти не індивідуальні показники, а індивідуальні ієрархії, або профілі, що недоступно жодному з інших статистичних методів, включно з методом лінійної кореляції. Коефіцієнт рангової кореляції рекомендується застосовувати в тих випадках, коли нам необхідно перевірити, чи

узгоджено змінюються різні ознаки в одного й того самого випробовуваного і наскільки збігаються індивідуальні рангові показники у двох окремих випробовуваних або у випробовуваного і групи

Ранжування

Часто у статистичному аналізі, наприклад, при порівняння сукупностей даних, використовують процедуру ранжування. Ця процедура призначає значенням, які вимірюються принаймні за порядковою шкалою, певна ранги. При цьому ранги пам'ятають порядок значень, але забувають відстань між значеннями. Опишемо цю процедуру. Нехай маємо деяку сукупність даних обсягом n . Впорядкуємо її за зростанням. Припустимо, одержано таку послідовність даних: $X_1, X_2, X_3, \dots, X_n$. Можливі два випадки: ідеальний, коли всі значення в сукупності різні, і такий, коли деякі значення повторюються.

Ідеальний випадок. У ситуації, коли значення в сукупності різні, найменшому значенню (елементу) X_1 , призначимо ранг 1, наступному за розміром елементу X_2 – ранг 2 і так до кінця. Найбільший елемент X_n матиме ранг n .

Випадок з повтореннями. У цій ситуації спочатку призначимо кожному елементу ранги так, ніби повторень немає. Назвемо ці ранги попередніми (табл. 5.1).

Таблиця 5.1

Впорядковані значення	X_1	X_2	X_3	...	X_{n-1}	X_n
Попередні ранги	1	2	3	...	$n-1$	n

Якщо деяке значення не має повторень, його попередній ранг i є рангом. Для однакових елементів попередні ранги необхідно модифікувати. Покажемо, як це зробити. Розглянемо деяку групу однакових значень (нехай решта даних у сукупності мають інші значення) з їх попередніми рангами (табл. 5.2).

Таблиця 5.2

Група однакових значень	X_i	X_{i+1}	...	X_{i+k}
Попередні ранги	i	$i+1$	...	$i+k$

Елементом цієї групи призначимо однаковий ранг, який дорівнює середньому значенню їх попередніх рангів:

$$\frac{i + (i + 1) + \dots + (i + k)}{k + 1}.$$

Зауважимо, що значення цього виразу можна обчислити простіше. Воно дорівнює середньому значенню двох крайніх елементів групи, тобто

$$\frac{i + (i + k)}{2}.$$

Процедура ранжування не складна, а скоріше механічна. Проте дослідники, виконуючи ранжування вручну, часто помиляються. Для контролю правильності ранжування прийнято перевіряти таку властивість: якщо сукупність має обсяг n , то загальна сума рангів її елементів дорівнює

$$\frac{n(n + 1)}{2}.$$

Приклад 5.1. Нехай психолог за допомогою тесту оцінив рівень роздратованості у групі вчителів початкових класів за семибальною шкалою. Наведемо результати дослідження (табл. 5.3).

Таблиця 5.3

№ вчит	1	2	3	4	5	6	7	8	9	10	11	12	13	14
Рівень роздрат.	1	5	2	5	3	3	7	1	2	2	2	4	5	1

Проранжуємо цей список. Для цього запишемо послідовність рівнів роздратованості в порядку зростання

1, 1, 1, 2, 2, 2, 2, 3, 3, 4, 5, 5, 5, 7.

Призначимо цим значенням попередні ранги (табл. 5.4).

Таблиця 5.4

Рівень роздрат.	1	1	1	2	2	2	2	3	3	4	5	5	5	7
Попередній ранг	1	2	3	4	5	6	7	8	9	10	11	12	13	14

Після цього обчислимо ранги в кожній підгрупі однакових значень. Для підгрупи елементів з рівнем 1 маємо попередні ранги 1, 2, 3. Тому ранг елементів підгрупи дорівнює

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

Попередні ранги підгрупи з рівнем роздратованості 2 дорівнюють 4, 5, 6 та 7. Крайнім позиціям відповідають ранги 4 і 7. Отже, спільний ранг підгрупи дорівнює

$$\frac{4 + 7}{2} = 5,5$$

Аналогічно знаходимо ранги підгруп з рівнями 3 і 5:

$$\frac{8 + 9}{2} = 8,5; \quad \frac{11 + 13}{2} = 12$$

Оскільки рівні 4 і 7 мають лише по одному індивідууму, їх ранги збігаються з попередніми рангами відповідно 10 і 14.

Запишемо результати ранжування (табл. 5.5).

Таблиця 5.5

Рівень роздратованості	1	2	3	4	5	7
Ранг	2	5,5	8,5	10	12	14

Перевіримо правильність ранжування, враховуючи кількості повторень рангів:

Сума рангів

$$2 \cdot 3 + 5,5 \cdot 4 + 8,5 \cdot 2 + 10 + 12 \cdot 3 + 14 = 105.$$

$$\frac{n(n+1)}{2} = \frac{14(14+1)}{2} = 7 \cdot 15 = 105.$$

Оскільки контрольна рівність правильна, маємо всі підстави вважати, що ранжування було здійснено без помилок. Остаточні результати ранжування зведемо в таблицю 5.6.

Таблиця 5.6

№ вчит	1	2	3	4	5	6	7	8	9	10	11	12	13	14
Рівень роздрат.	1	5	2	5	3	3	7	1	2	2	2	4	5	1
Ранг	2	12	5,5	12	8,5	8,5	14	2	5,5	5,5	5,5	10	12	2

5.2 Коефіцієнт рангової кореляції r_s Спірмена

Метод рангової кореляції Спірмена дає змогу визначити тісноту (силу) і напрямок кореляційного зв'язку між двома ознаками або двома профілями (ієрархіями) ознак.

Для підрахунку рангової кореляції необхідно мати у своєму розпорядженні два ряди значень, які можуть бути проранжовані. Такими рядами значень можуть бути:

1) дві ознаки, виміряні в одній і тій самій групі випробовуваних;

2) дві індивідуальні ієрархії ознак, виявлені у двох випробовуваних за одним і тим самим набором ознак;

3) дві групові ієрархії ознак,

4) індивідуальна і групова ієрархії ознак.

Спочатку показники ранжуються окремо за кожною з ознак.

Як правило, меншому значенню ознаки нараховується менший ранг.

У першому випадку (дві ознаки) ранжуються індивідуальні значення за першою ознакою, отримані різними випробовуваними, а потім індивідуальні значення за другою ознакою.

Якщо дві ознаки пов'язані позитивно, то випробовувані, які мають низькі ранги за однією з них, матимуть низькі ранги і за іншою, а випробовувані, які мають високі ранги за однією з ознак, матимуть за іншою ознакою також високі ранги. Для підрахунку r_s необхідно визначити різниці (d) між рангами, отриманими даним випробовуваним за обома ознаками. Потім ці показники d певним чином перетворюються і віднімаються від 1. Чим менша різниця між рангами, тим більшим буде r_s , тим ближче він буде до +1.

Якщо кореляція відсутня, то всі ранги будуть перемішані і між ними не буде жодної відповідності. Формула складена так, що в цьому випадку r_s виявиться близьким до 0.

У разі негативної кореляції низьким рангам випробовуваних за однією ознакою відповідатимуть високі ранги за іншою ознакою, і навпаки. Чим більша розбіжність між рангами випробовуваних за двома змінними, тим ближче r_s до -1.

У другому випадку (два індивідуальні профілі), ранжуються індивідуальні значення, отримані кожним із 2-х випробовуваних за певним (однаковим для них обох) набором ознак. Перший ранг отримає ознака з найнижчим значенням; другий ранг – ознака з вищим значенням і т.д. Очевидно, що всі ознаки мають бути виміряні в одних і тих самих одиницях, інакше ранжування неможливе.

Якщо індивідуальні ієрархії двох випробовуваних пов'язані позитивно, то ознаки, що мають низькі ранги в одного з них, матимуть низькі ранги й в іншого, і навпаки. Наприклад, якщо в одного випробовуваного фактор Е (домінантність) має найнижчий ранг, то й в іншого випробовуваного він повинен мати низький ранг, якщо в одного випробовуваного фактор С (емоційна стійкість) має вищий ранг, то й інший випробовуваний повинен мати за цим фактором високий ранг і т. п.

У третьому випадку (два групових профілі), ранжуються середньогрупові значення, отримані в 2-х групах випробовуваних за певним, однаковим для двох груп, набором ознак. Надалі лінія міркувань така сама, як і в попередніх двох випадках.

У випадку 4-му (індивідуальний і груповий профілі), ранжуються окремо індивідуальні значення випробовуваного і середньогрупові значення за тим самим набором ознак, які отримані, як правило, за умови виключення цього окремого випробовуваного – він не бере участі в середньогруповому профілі, з яким буде зіставлятися його індивідуальний профіль. Рангова кореляція дасть змогу перевірити, наскільки узгоджені індивідуальний і груповий профілі.

У всіх чотирьох випадках значущість отриманого коефіцієнта кореляції визначається за кількістю ранжованих значень N . У першому випадку ця кількість збігатиметься з обсягом вибірки n . У другому випадку кількістю спостережень буде кількість ознак, що складають ієрархію. У третьому і четвертому випадку N – це також кількість зіставляваних ознак, а не кількість випробовуваних у групах. Детальні пояснення наведено в прикладах. Якщо абсолютна величина r_s досягає критичного значення або перевищує його, кореляція достовірна.

Гіпотези.

Можливі два варіанти гіпотез. Перший стосується випадку 1, другий – до трьох інших випадків.

Перший варіант гіпотез

H_0 : Кореляція між змінними A і B не відрізняється від нуля.

H_1 : Кореляція між змінними A і B достовірно відрізняється від нуля.

Другий варіант гіпотез

H_0 : Кореляція між ієрархіями A і B не відрізняється від нуля.

H_1 : Кореляція між ієрархіями A і B достовірно відрізняється від нуля.

Обмеження коефіцієнта рангової кореляції

1. За кожною змінною має бути представлено не менше 5 спостережень. Верхня межа вибірки визначається наявними таблицями критичних значень.

2. Коефіцієнт рангової кореляції Спірмена r_s за великої кількості однакових рангів за однією або обома зіставляваними змінними дає не точні значення. В ідеалі обидва корельовані ряди мають являти собою дві послідовності значень, що не збігаються. У разі, якщо ця умова не дотримується, необхідно вносити поправку на однакові ранги.

Коефіцієнт рангової кореляції Спірмена підраховується за формулою:

$$r_s = 1 - \frac{6 \sum d^2}{N(N^2 - 1)},$$

де d – різниця між рангами за двома змінними для кожного випробовуваного; N – кількість ранжованих значень, в даному випадку кількість випробовуваних.

Якщо в обох зіставляваних рангових рядах присутні групи однакових рангів, перед підрахунком коефіцієнта рангової кореляції необхідно внести поправки на однакові ранги:

$$T_a = \sum \frac{a^3 - a}{12};$$

$$T_b = \sum \frac{b^3 - b}{12};$$

де a – обсяг кожної групи однакових рангів у ранговому ряду A ,
 b – обсяг кожної групи однакових рангів у ранговому ряду B .

Для підрахунку емпіричного значення r_s використовують формулу:

$$r_s = 1 - 6 \cdot \frac{\sum d^2 + T_a + T_b}{N(N^2 - 1)},$$

Розрахунок коефіцієнта рангової кореляції Спірмена r_s

1. Визначити, які дві ознаки або дві ієрархії ознак братимуть участь у зіставленні як змінні A і B .

2. Проранжувати значення змінної A , нараховуючи ранг 1 найменшому значенню, відповідно до правил ранжування. Занести ранги в перший стовпець таблиці за порядком номерів випробовуваних або ознак.

3. Проранжувати значення змінної B , відповідно до тих самих правил. Занести ранги в другий стовпець таблиці за порядком номерів випробовуваних або ознак.

4. Підрахувати різниці d між рангами A і B за кожним рядком таблиці і занести в третій стовпчик таблиці.

5. Підвести кожен різницю до квадрата: d^2 . Ці значення занести в четвертий стовпець таблиці.

6. Підрахувати суму квадратів $\sum d^2$.

7. За наявності однакових рангів розрахувати поправки:

$$T_a = \sum \frac{a^3 - a}{12};$$
$$T_b = \sum \frac{b^3 - b}{12};$$

де a – обсяг кожної групи однакових рангів у ранговому ряду A ,
 b – обсяг кожної

групи однакових рангів у ранговому ряду B .

8. Розрахувати коефіцієнт рангової кореляції r_s за формулою:

а) за відсутності однакових рангів

$$r_s = 1 - \frac{6 \sum d^2}{N(N^2 - 1)},$$

б) за наявності однакових рангів

$$r_s = 1 - 6 \cdot \frac{\sum d^2 + T_a + T_b}{N(N^2 - 1)},$$

де $\sum d^2$ – сума квадратів різниць між рангами, T_a і T_b – поправки на однакові ранги;

N – кількість випробовуваних або ознак, які брали участь у ранжуванні.

9. Визначити за таблицею критичні значення r_s для даного N . Якщо r_s , перевищує критичне значення або принаймні дорівнює йому, кореляція достовірно відрізняється від 0.

Приклад 5.2 При визначенні ступеня залежності реакції вживання алкоголю на окорухову реакцію у випробовуваній групі було отримано дані до вживання алкоголю і після вживання. Чи залежить реакція випробуваного від стану сп'яніння?

Результати експерименту:

До: 16, 13, 14, 9, 10, 13, 14, 14, 18, 20, 15, 10, 9, 10, 16, 17, 18.

Після: 24, 6, 9, 10, 23, 20, 11, 12, 19, 18, 13, 14, 12, 14, 7, 9, 14.

Сформулюємо гіпотези:

H_0 : кореляція між ступенем залежності реакції до вживання алкоголю і після не відрізняється від нуля.

H_1 : кореляція між ступенем залежності реакції до вживання алкоголю та після достовірно відрізняється від нуля.

Встановимо ранги до експерименту (табл. 5.7):

Таблиця 5.7

рівень	9	9	10	10	10	13	13	14	14	14	15	16	16	17	18	18	20
п. ранг	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17
ранг	1,5	1,5	4	4	4	6,5	6,5	9	9	9	11	12,5	12,5	14	15,5	15,5	17

Встановимо ранги після експерименту (табл. 5.8):

Таблиця 5.8

рівень	6	7	9	9	10	11	12	12	13	14	14	14	18	19	20	23	24
п. ранг	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17
ранг	1	2	3,5	3,5	5	6	7,5	7,5	9	11	11	11	13	14	15	16	17

Таблиця 5.9. Розрахунок d^2 для рангового коефіцієнта кореляції Спірмена r_s при зіставлення показників окорухової реакції до експерименту і після ($N = 17$).

№ п/п	До		Після		d	d^2
	значення	ранг	значення	ранг		
1	2	3	4	5	6	7
1	16	12,5	24	17	-4,5	20,25
2	13	6,5	6	1	5,5	30,25
3	14	9	9	3,5	5,5	30,25
4	9	1,5	10	5	-3,5	12,25
5	10	4	23	16	-12	144
6	13	6,5	20	15	-8,5	72,25
7	14	9	11	6	3	9
8	14	9	12	7,5	1,5	2,25
9	18	15,5	19	14	1,5	2,25
10	20	17	18	13	4	16
11	15	11	13	9	2	4

Продовження таблиці 5.9

1	2	3	4	5	6	7
12	10	4	14	11	-7	49
13	9	1,5	12	7,5	-6	36
14	10	4	14	11	-7	49
15	16	12,5	7	2	10,5	110,25
16	17	14	9	3,5	10,5	110,25
17	18	16	14	11	5	25

$$\Sigma d^2 = 722,25.$$

Оскільки ми маємо ранги, що повторюються, то в даному випадку будемо застосовувати формулу з поправкою на однакові ранги:

$$r_s = 1 - 6 \cdot \frac{\Sigma d^2 + T_a + T_b}{N(N^2 - 1)},$$

$$T_a = \sum \frac{a^3 - a}{12};$$

$$T_a = \frac{1}{12} ((2^3 - 2) + (3^3 - 3) + (2^3 - 2) + (3^3 - 3) + (2^3 - 2) + (2^3 - 2)) =$$

$$= \frac{1}{12} (6 + 24 + 6 + 24 + 6 + 6) = \frac{72}{12} = 6;$$

$$T_b = \sum \frac{b^3 - b}{12};$$

$$T_b = \frac{1}{12} ((2^3 - 2) + (2^3 - 2) + (3^3 - 3)) = \frac{1}{12} (6 + 6 + 24) = 3.$$

$$r_s = 1 - 6 \cdot \frac{722,25 + 6 + 3}{17 \cdot (17^2 - 1)} = 0,1039$$

За таблицею критичних значень, знаходимо критичні значення коефіцієнта кореляції для N=17:

$$r_{кр} = \begin{cases} 0,48 & (p \leq 0,05) \\ 0,62 & (p \leq 0,01) \end{cases} \Rightarrow r_s < r_{кр}$$

Висновок: H_1 гіпотеза відкидається і приймається H_0 . Тобто кореляція між ступенем залежності реакції до вживання алкоголю і після не відрізняється від нуля.

5.3 Коефіцієнт кореляції Браве-Пірсона

Для обчислення цього коефіцієнта застосовують таку формулу (у різних авторів вона може виглядати по-різному):

$$r = \frac{(\sum x_i y_i) - n \cdot \bar{M}_1 \cdot \bar{M}_2}{(n - 1) \cdot \sigma_x \cdot \sigma_y},$$

де $\sum x_i y_i$ – сума добутків даних з кожної пари, n – число пар, $\bar{M}_1$ – середнє для даних змінної X , $\bar{M}_2$ – середнє для даних змінної Y , σ_x – стандартне відхилення для розподілу x , σ_y – стандартне відхилення для розподілу y .

Оскільки в прикладі 5.2 з попереднього підрозділу розподіл обох ознак не відрізняється від нормального, то ми можемо знайти коефіцієнт лінійної кореляції Браве-Пірсона і встановити, чи існує залежність між вживанням алкоголю і реакцією випробуваного (табл. 5.10).

Таблиця 5.10

№ п/п	Значення		$x_i y_i$
	до(x_i)	після(y_i)	
1	16	24	384
2	13	6	78
3	14	9	126
4	9	10	90
5	10	23	230
6	13	20	260
7	14	11	154
8	14	12	168
9	18	19	342
10	20	18	360
11	15	13	195
12	10	14	140
13	9	12	108
14	10	14	140
15	16	7	112
16	17	9	153
17	18	14	252

$$\sum x_i y_i = 3292$$

$$\bar{M}_1 = 13,9; \bar{M}_2 = 13,8; \sigma_x = 3,4; \sigma_y = 5,3$$

$$r = \frac{3292 - 17 \cdot 13,9 \cdot 13,8}{16 \cdot 3,4 \cdot 5,3} = \frac{31,06}{290,6} = 0,11.$$

Критичні значення коефіцієнта кореляції такі ж, як і в попередньому прикладі.

$$r_{\text{кр}} = \begin{cases} 0,48 & (p \leq 0,05) \\ 0,62 & (p \leq 0,01) \end{cases} \Rightarrow r_{\text{емп}} < r_{\text{кр}}$$

Висновок: H_1 гіпотеза відкидається і приймається H_0 . Тобто кореляція між ступенем залежності реакції до вживання алкоголю і після не відрізняється від нуля.

5.4 Інтерпретація коефіцієнтів кореляції

Причинність і кореляція. Наявність кореляції двох змінних аж ніяк не означає, що між ними існує причинний зв'язок. Незважаючи на те, що співіснування (кореляцію) подій можна використовувати для виявлення причинних зв'язків поряд з іншими методологічними підходами, монопольне застосування кореляції до аналізу причинності ризиковане і може вводити в оману.

По-перше, навіть у тих випадках, коли можна припустити існування причинного зв'язку між двома змінними, які корельовані, r сам по собі нічого не говорить про те, чи викликає x появу y , чи y викликає появу x .

По-друге, часто спостережуваний зв'язок існує завдяки іншим змінним, а не двом розглядуваним.

По-третє, взаємозв'язки змінних у педагогіці та суспільних науках майже завжди занадто складні, щоб їх поясненням могла слугувати єдина причина.

Багато дослідників не зупиняються на тому хибному висновку, що кореляція свідчить *на перший погляд* про причинну залежність, а виводять також і інший висновок. Вони приписують причинному зв'язку певний напрямок.

Розглянемо більш уважно правдоподібний приклад.

Хоча кореляція прямо не вказує на причинний зв'язок, вона може слугувати ключем до розгадки причин. За сприятливих умов на її основі можна сформулювати гіпотези, які перевіряють експериментально, коли можливий контроль інших впливів, крім тих нечисленних, які підлягають дослідженню. Існують також добре розроблені процедури, зокрема в соціології та психології, для виведення причин із пов'язаних даних.

Іноді відсутність кореляції може мати більш глибокий вплив на нашу гіпотезу про причинний зв'язок, ніж наявність сильної кореляції. Нульова кореляція двох змінних може свідчити про те, що жодного впливу однієї змінної на іншу не існує, за умови, що ми довіряємо результатам вимірювань і що добуток моментів r Пірсона, що вимірює тільки приватний тип зв'язку, підходить для вимірювання більш загального типу зв'язку, званого «причинним». Але все це мало допомагає: потрібні методи виявлення причинних зв'язків, а не методи ілюстрації безпричинних явищ.

Ідентичні групи з різними середніми. Істотна кореляція між двома змінними – це факт, який у різних ситуаціях можна пояснити по-різному. Деякі кореляції – результат вимірювання причини та її дії, наприклад, коли X – їжа, з'їдена за місяць, а Y – вага, набута за той самий час. Інші кореляції виникають при вимірюваннях двох змінних зі спільною причиною або впливом, наприклад коли X – успішність з англійської мови, а Y – із суспільних наук. Іноді виникають інші кореляції, коли об'єднуються дві різні групи, у кожній з яких X і Y не мають зв'язку.

ЛЕКЦІЯ 6

МЕТОДИ ПЕРЕВІРКИ СТАТИСТИЧНИХ ГІПОТЕЗ

1. t -критерій Стьюдента.
2. F -критерій Фішера (для порівняння дисперсій).
3. Виявлення відмінностей у рівні досліджуваної ознаки. Q -критерій Розенбаума.
4. T -критерій Вілкоксона.
5. Виявлення відмінностей у розподілі ознаки. χ^2 -критерій Пірсона.

6.1 t -критерій Стьюдента

t -критерій Стьюдента використовується для:

1) встановлення подібності-відмінності середніх арифметичних значень у двох вибірках ($M_1 \leftrightarrow M_2$) або, у більш загальному вигляді, для встановлення подібності-відмінності двох емпіричних розподілів;

2) встановлення відмінності від нуля деяких заходів зв'язку: коефіцієнта лінійної кореляції Пірсона, рангової кореляції Спірмена;

3) встановлення схожості-відмінності двох дисперсій у двох залежних вибірках.

Обмеження:

1) це параметричний критерій, тому необхідно, щоб розподіл ознаки, принаймні не відрізнявся від нормального розподілу;

2) для незалежних і залежних вибірок різні формули розрахунку;

Гіпотези

1) незалежні вибірки:

H_0 : середні значення ознаки в обох вибірках не відрізняються,

H_1 : середні значення ознаки в обох вибірках статистично значуще різняться.

2) залежні вибірки:

H_0 : різниці оцінок випробовуваних у двох станах не відрізняються від нуля,

H_1 : різниці оцінок випробовуваних у двох станах статистично значимо відрізняються від нуля.

Розглянемо випадок 1.

Приклад 6.1 (незалежні вибірки). Припустимо, є дві незалежні вибірки школярів, інтелект яких розвивали протягом деякого часу за двома різними методиками, необхідно встановити, яка з методик краща (табл. 6.1). Попередньо було з'ясовано, що початковий рівень інтелекту був однаковим в обох вибірках. Завдання порівняння двох методик може бути переформульоване мовою статистики як завдання порівняння середніх арифметичних значень інтелекту в обох вибірках .

Таблиця 6.1

<i>Числові характеристики</i>	<i>1-ша вбірка</i>	<i>2-га вбірка</i>
<i>n</i>	30	32
<i>M</i>	103	110
<i>σ</i>	10	12

Гіпотези:

H_0 : середні значення рівня інтелекту в обох вибірках не відрізняються,

H_1 : середні значення рівня інтелекту в обох вибірках статистично значимо розрізняються.

У цьому випадку для отримання емпіричного значення t-критерію використовується така формула:

$$t = \frac{|M_1 - M_2|}{\sqrt{(n_1 - 1)\sigma_1^2 + (n_2 - 1)\sigma_2^2}} \sqrt{\frac{n_1 \cdot n_2 \cdot (n_1 + n_2 - 2)}{n_1 + n_2}} .$$

де: n_1, n_2 – кількість випробовуваних у 1-й і 2-й вибірках; M_1, M_2 – середні арифметичні значення в 1-й і 2-й вибірках; σ_1, σ_2 – стандартні відхилення в 1-й і 2-й вибірках.

Кількість ступенів свободи для знаходження критичного значення критерію:

$$Df = n_1 + n_2 - 2$$

(У розглянутому прикладі критичні значення t -критерію наводяться для ненаправлених гіпотез).

Тоді:

$$t = \frac{|103 - 110|}{\sqrt{29 \cdot 100 + 31 \cdot 144}} \sqrt{\frac{30 \cdot 32 \cdot (30 + 32 - 2)}{30 + 32}} \approx 2,486.$$

Таким чином, отримуємо $t_{\text{емп}} \approx 2,486$

Критичні значення t -критерію знаходимо за таблицею для

$$Df = 30 + 32 - 2 = 60$$

$$t_{\text{кр}} = \begin{cases} 2,0 & \text{для } p \leq 0,05 \\ 2,66 & \text{для } p \leq 0,01 \end{cases}$$

Отримане емпіричне значення t -критерію перевищує критичне для $\alpha = 0,05$, але виявляється меншим за критичне для $\alpha = 0,01$, тобто

$$2,0 < t_{\text{емп}} \approx 2,486 < 2,66$$

Висновок: H_0 гіпотеза відхиляється і можна зробити висновок про статистично значущу відмінності середніх арифметичних значень у двох вибірках для $p \leq 0,05$ і про переваги другої методики порівняно з першою.

Суворе використання t -критерію передбачає, що обидві вибірки витягнуті з нормальних сукупностей. Однак багато авторів не вважають цю умову достатньо жорсткою, вказуючи на можливість використання t -критерію в ситуаціях, коли немає серйозних підстав сумніватися в нормальності розподілу ознаки в генеральній сукупності, навіть якщо це не можна підтвердити статистично.

У разі залежних вибірок виникає кореляція результатів, оскільки вимірювання проводять на одних і тих самих

випробовуваних у різних умовах $(x \text{ і } y)'$, щоб урахувати вплив кореляції, застосовується інша формула:

$$t = \frac{\sum d_i}{\sqrt{\sum d_i^2 - (\sum d_i)^2/n}} \sqrt{\frac{n-1}{n}}.$$

де $d_i = x_i - y_i$, тобто різниця значень ознаки для кожного випробовуваного. Кількість ступенів свободи $df = n - 1$. Перевіряється статистична гіпотеза про відповідність розподілу різниць t -розподілу Стюдента з нульовим середнім значенням.

Приклад 6.2. (залежні вибірки). Припустимо, проводиться вимірювання ситуативної тривожності до і після психотерапевтичного впливу за допомогою деякого опитувальника (табл. 6.2). Дослідника цікавить питання, чи призводить вплив до зміни рівня тривожності.

Гіпотези:

H_0 : різниці оцінок у випробовуваних ситуативної тривожності до і після психотерапевтичного впливу не відрізняються від нуля,

H_1 : різниці оцінок у випробовуваних ситуативної тривожності до і після психотерапевтичного впливу статистично значимо відрізняються від нуля.

Таблиця 6.2

Випробовувані	«до» x_i	«після» y_i	d_i	d_i^2
1	30	20	10	100
2	33	17	16	256
3	41	21	20	400
4	50	43	7	49
5	36	39	-3	9
6	45	11	34	1156
7	31	28	3	9
8	25	20	5	25
$n = 8$			$\sum d_i$ $= 92$	$\sum d_i^2$ $= 2004$

$$t = \frac{92}{\sqrt{2004 - \frac{92^2}{8}}} \cdot \sqrt{\frac{7}{8}} \approx 2,798$$

Таким чином, отримуємо $t_{\text{емп}} \approx 2,798$

Критичні значення t -критерію знаходимо за таблицею для

$$Df = 8 - 1 = 7$$

$$t_{\text{кр}} = \begin{cases} 2,365 & \text{для } p \leq 0,05 \\ 3,499 & \text{для } p \leq 0,01 \end{cases}$$

Отримане емпіричне значення t -критерію перевищує критичне для $\alpha = 0,05$, але виявляється меншим за критичне для $\alpha = 0,01$, тобто

$$2,365 < t_{\text{емп}} \approx 2,798 < 3,499$$

Висновок: Приймається H_1 гіпотеза. Відмінності в рівнях тривожності до і після психотерапевтичного впливу слід визнати статистично значущими ($p < 0,05$), оскільки емпіричне значення перевищує перше критичне, але менше другого. Отже, психотерапевтичний вплив дійсно знижує тривожність.

Випадок 2. Під час перевірки відмінності від нуля мір зв'язку (коефіцієнтів кореляції) емпіричне значення t -критерію обчислюється за формулою

$$t = r \sqrt{\frac{n-2}{1-r^2}}$$

де r – коефіцієнт кореляції, n – кількість випробовуваних. Кількість ступенів свободи $df = n-2$. Висновок про відмінність міри зв'язку від нуля робиться при перевищенні емпіричного значення критерію над критичним для $\alpha = 0,05$ і відповідної кількості ступенів свободи, тобто аналогічно розглянутим вище випадкам.

Для коефіцієнтів лінійної кореляції Пірсона і рангової кореляції Спірмена

можна безпосередньо використовувати таблиці критичних значень. Емпіричне значення коефіцієнта кореляції береться за абсолютною величиною.

У деяких посібниках і підручниках наводяться окремі таблиці критичних значень для коефіцієнта рангової кореляції Спірмена, за В. Ю. Урбахом. Значення в них відрізняються від критичних для коефіцієнтів лінійної кореляції Пірсона.

6.2 *F*-критерій Фішера (для порівняння дисперсій)

F-критерій Фішера використовується для:

1) встановлення подібності-відмінності дисперсій у двох незалежних вибірках ($D_1 \leftrightarrow D_2$);

2) встановлення відмінності від нуля коефіцієнта детермінації ($\eta^2 \leftrightarrow \gg 0$);

3) встановлення наявності-відсутності впливу фактора в дисперсійному аналізі.

Випадок 1

Емпіричне значення *F*-критерію для порівняння двох дисперсій у незалежних вибірках знаходять за дуже простою формулою:

$$F = \frac{D_1}{D_2} = \frac{\sigma_1^2}{\sigma_2^2}$$

де D_1 – більша дисперсія, D_2 – менша дисперсія (Підстановка в чисельник більшої дисперсії необхідна для використання таблиць критичних значень, у яких наводиться тільки праве критичне значення (більше одиниці). Статистичні програми розраховують і ліве критичне значення (менше одиниці)).

Кількість ступенів свободи визначається окремо для чисельника й окремо для знаменника:

$$df_{\text{чисел.}} = n_{\text{чисел.}} - 1$$

$$df_{\text{знам.}} = n_{\text{знам.}} - 1$$

Приклад 6.3. Дві групи випробовуваних навчалися деяких моторних навичок за двома різними методиками, фіксувалася кількість помилкових дій, до навчання результати в обох групах мали однаковий розкид (табл. 6.3). Яка з методик дасть найбільше вирівнювання результатів усередині групи після навчання.

Таблиця 6.3

Числові характеристики	1-ша група	2-га група
n	21	16
D	16	36
σ	4	4

$$F = \frac{D_1}{D_2} = \frac{36}{16} = 2,25;$$

$$df_{\text{чисел.}} = n_{\text{чисел.}} - 1 = 15;$$

$$df_{\text{знам.}} = n_{\text{знам.}} - 1 = 20.$$

Оскільки нам заздалегідь не відомо, яка з методик може мати меншу дисперсію, ми використовуємо ненаправлену гіпотезу і, отже, двосторонній критерій. Знаходимо за таблицею критичне значення $F_{\text{кр.}}$. Для $\alpha = 0,05$ ($\alpha/2 + \alpha/2 = 0,05$) і $df_{\text{чисел.}} = 15$; $df_{\text{знам.}} = 20$, $F_{\text{кр.}} = 2,573$.

Отримаємо:

$$F_{\text{емп.}} = 2,25 < F_{\text{кр.}} = 2,573$$

Висновок: Оскільки емпіричне значення менше критичного, то статистично значущих відмінностей дисперсій у першій і другій групах немає і, отже, стабілізація навички при навчанні за обома методиками однакова.

Зауваження. Для порівняння дисперсій у залежних вибірках суворішим буде застосування t-критерію Стюдента.

6.3 Виявлення відмінностей у рівні досліджуваної ознаки. Q-критерій Розенбаума

У психології часто доводиться проводити *дослідження на виявлення відмінностей між двома, трьома і більше вибірками випробовуваних*. Це може бути, наприклад, завдання визначення психологічних особливостей хворих дітей порівняно зі здоровими, або відмінностей між працівниками державних підприємств і приватних фірм тощо.

Іноді за виявленими в дослідженні статистично достовірними відмінностями формується «груповий профіль» або «усереднений портрет» людини тієї чи іншої професії, статусу, соматичного захворювання тощо.

Останніми роками дедалі частіше постає завдання виявлення психологічного портрета фахівця нових професій: «успішного менеджера», «успішного політика», «успішного торгового представника», «успішного комерційного директора» тощо. Такого роду дослідження не завжди передбачають участь двох або більше вибірок. *Іноді обстежують одну, але доволі представницьку вибірку чисельністю не менш як 60 осіб, а потім усередині цієї вибірки виділяють групи більш і менш успішних фахівців*, і їхні дані за дослідженими змінними зіставляються між собою. У найпростішому випадку критерієм для поділу вибірки на «успішних» і «неуспішних» буде середня величина за показником успішності. Однак такий поділ є доволі грубим: особи, які отримали близькі оцінки за успішністю, можуть опинитися в протилежних групах, а особи, які помітно різняться за оцінками успішності, - в одній і тій самій групі.

Це може спотворити результати зіставлення груп або, принаймні, зробити відмінності між групами менш помітними.

Щоб уникнути цього, можна спробувати виділити групи «успішних» і «неуспішних» фахівців суворіше, включаючи в першу з них тільки тих, чиї значення перевищують середню величину не менш ніж на $1/4$ стандартного відхилення, а в другу групу – тільки тих, чиї значення не менш ніж на $1/4$ стандартного

відхилення нижчі за середню величину. При цьому всі, хто опиняється в зоні середніх величин, $M \pm 1/4 \sigma$, випадають із подальших зіставлень. Якщо розподіл близький до нормального, то випаде приблизно 19,8% випробовуваних. Якщо розподіл відрізняється від нормального, то таких випробовуваних може бути й більше. Щоб уникнути втрат, можна зіставляти не дві, а три групи випробовуваних: з високою, середньою і низькою професійною успішністю.

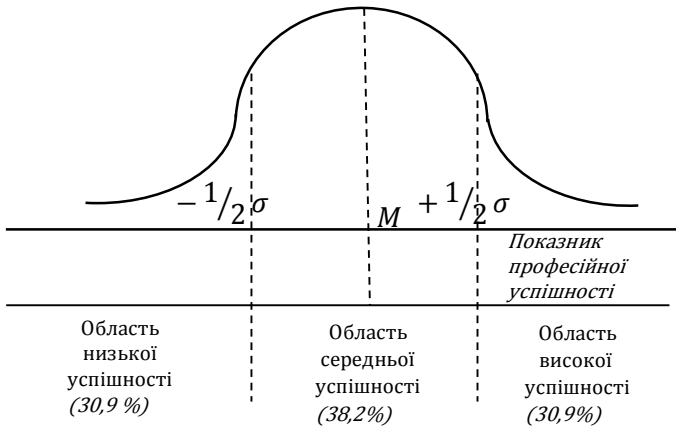


Рисунок 6.1.

На рисунку 6.1 представлено схему поділу вибірки на групи з низькою, середньою та високою професійною успішністю за критерієм відхилення значень від середньої величини на $1/2$ стандартного відхилення. За такого суворого критерію в «середню» групу потрапляють (при нормального розподілу) близько 38,2% усіх випробовуваних, а в крайніх групах виявляється по 30,9% випробовуваних.

Чим менше випробовуваних опиняється в групах, тим менше у нас можливостей для виявлення достовірних відмінностей, оскільки критичні значення більшості критеріїв за малих n суворіші, ніж за великих n .

Таким чином, у разі не суворого поділу випробовуваних на групи ми втрачаємо в точності, а за суворого – у кількості випробовуваних.

Зіставлення рівневих показників у різних вибірках може бути необхідною частиною комплексних діагностичних, навчальних, психокорекційних та інших програм. Воно допомагає звернути увагу на ті особливості обстежених вибірок, які мають бути враховані й використані під час адаптації програм до даної групи в процесі їх конкретного втілення.

Рішення про вибір того чи іншого критерію ухвалюють на основі того, скільки вибірок зіставляється і який їхній обсяг.

Призначення критерію.

Критерій використовують для оцінювання відмінностей між двома вибірками за рівнем будь-якої ознаки, кількісно вимірної. У кожній із вибірок має бути не менше 11 випробовуваних.

Опис критерію.

Це дуже простий *непараметричний критерій*, який дає змогу швидко оцінити відмінності між двома вибірками за будь-якою ознакою. Однак якщо критерій Q не виявляє достовірних відмінностей, це ще не означає, що їх справді немає.

У цьому разі варто застосувати критерій ϕ – Фішера. Якщо ж Q -критерій виявляє достовірні відмінності між вибірками з рівнем значущості $p \leq 0,01$, то можна обмежитися тільки ним і уникнути труднощів застосування інших критеріїв.

Критерій застосовують у тих випадках, коли дані подано принаймні в порядковій шкалі. Ознака має варіювати в якомусь діапазоні значень, інакше зіставлення за допомогою Q -критерію просто неможливі. Наприклад, якщо у нас тільки 3 значення ознаки, 1, 2 і 3 – нам дуже важко буде встановити відмінності. Метод Розенбаума вимагає досить тонко вимірних ознак.

Застосування критерію починається з упорядкування значень ознаки в обох вибірках за наростанням (або спаданням) ознаки. Для того щоб не заплутатися, у цьому та в багатьох інших

критеріях рекомендується першим рядом (вибіркою, групою) вважати той ряд, де значення вищими, а другим рядом – той, де значення нижчі.

Гіпотези.

H_0 : Рівень ознаки у вибірці 1 не перевищує рівня ознаки у вибірці 2.

H_1 : Рівень ознаки у вибірці 1 перевищує рівень ознаки у вибірці 2.

Обмеження критерію Q

1) У кожній із зіставляваних вибірок має бути не менше 11 спостережень. При цьому обсяги вибірок мають приблизно збігатися:

а) якщо в обох вибірках менше 50 спостережень, то абсолютна величина різниці між n_1 і n_2 не повинна бути більшою за 10 спостережень;

б) якщо в кожній із вибірок більше 51 спостереження, але менше 100, то абсолютна величина різниці між n_1 і n_2 не повинна бути більшою за 20 спостережень;

в) якщо в кожній із вибірок більше 100 спостережень, то допускається, щоб одна з вибірок була більшою за іншу не більше ніж у 1,5-2 рази.

2) Діапазони розкиду значень у двох вибірках повинні не збігатися між собою.

Приклад 6.4. У передбачуваних учасників психологічного експерименту, що моделює діяльність повітряного диспетчера, було виміряно рівень вербального і невербального інтелекту за допомогою методики Д. Векслера. Було обстежено 26 юнаків віком від 18 до 24 років (середній вік 20,5 років). 14 із них були студентами фізичного факультету, а 12 – студентами психологічного факультету Київського університету (табл. 6.4).

Чи можна стверджувати, що одна з груп перевершує іншу за рівнем вербального інтелекту?

Таблиця 6.4

<i>Студенти-фізики</i>			<i>Студенти-психологи</i>		
<i>Код імені випробуваного</i>		<i>Показники вербальн. Інтелекту</i>	<i>Код імені випробуваного</i>		<i>Показники вербальн. Інтелекту</i>
1	I.A.	132	1	H.T.	126
2	K.O.	134	2	O.B.	127
3	K.C.	124	3	C.M.	132
4	P.M.	132	4	F.O.	120
5	C.O.	135	5	I.M.	119
6	S.A.	132	6	I.K.	126
7	T.O.	131	7	L.D.	120
8	F.D.	132	8	K.I.	123
9	C.I.	121	9	D.Y.	120
10	C.O.	127	10	B.P.	116
11	S.M.D.	136	11	P.P.	123
12	K.A.	129	12	M.M.	115
13	B.L.	136			
14	F.V.	136			

Упорядкуємо значення в обох вибірках, а потім сформулюємо гіпотези:

H_0 : Студенти-фізики не перевершують студентів-психологів за рівнем вербального інтелекту.

H_1 : Студенти-фізики перевершують студентів-психологів за рівнем вербального інтелекту.

Як видно з таблиці 6,5, ми позначили ряди: перший, той, що «вище» – ряд фізиків, а другий, той, що «нижчий» – ряд психологів.

За таблиці 6.5 визначаємо кількість значень першого ряду, які більші за максимального значення другого ряду: $S_1=5$.

Тепер визначаємо кількість значень другого ряду, які менші за мінімальне значення першого ряду: $S_2=6$.

Як видно з таблиці 6.5 ми позначили ряди: перший, той, що «вище» – ряд фізиків, а другий, той, що «нижчий» – ряд психологів.

За таблицею 6.5 визначасмо кількість значень першого ряду, які більші за максимального значення другого ряду: $S_1=5$.

Таблиця 6.5

1-й ряд студенти-фізики		2-й ряд студенти-психологи	
1	См.Д. 136		
2	Б.Л. 136		
3	Ф.В. 136		
4	С.О. 135		
5	К.О. 134		
6	І.А. 132	1	Є.М. 132
7	П.М. 132		
8	Ст.А. 132		
9	Ф.Д. 132		
10	Т.О. 131		
11	К.Ан. 129		
12	Ц.О. 127	2	О.В. 127
		3	Н.Т. 126
		4	І.К. 126
13	К.Є. 124		
		5	К.І. 123
		6	Р.П. 123
14	Ч.І. 121		
		7	Ф.О. 120
		8	Л.Д. 120
		9	Д.Ю. 120
		10	І.М. 119
		11	В.Р. 116
		12	М.М. 115

Тепер визначасмо кількість значень другого ряду, які менші за мінімальне значення першого ряду: $S_2=6$.

Обчислюємо $Q_{\text{емп.}}$, за формулою:

$$Q_{\text{емп.}} = S_1 + S_2 = 11$$

За таблицею критичних значень визначасмо $Q_{\text{кр.}}$. Для вибірок об'ємом $n_1=14$ та $n_2=12$.

$$Q_{\text{кр.}} = \begin{cases} 7, & (p \leq 0,05) \\ 9, & (p \leq 0,01) \end{cases}$$

Чим більша розбіжність між вибірками, тим більша величина Q .

$$Q_{\text{кр.}} < Q_{\text{емп.}} \quad (p \leq 0,01)$$

Відповідь: гіпотеза H_0 відхиляється. Приймається гіпотеза H_1 . *Студенти-фізики перевершують студентів-психологів за рівнем вербального інтелекту ($p \leq 0,01$).*

Зауваження. У тих випадках, коли емпірична величина критерію опиняється на межі зони незначимості, ми маємо право стверджувати лише, що відмінності достовірні за $p \leq 0,05$, якщо ж воно виявляється між двома критичними значеннями, то ми можемо стверджувати, що $p < 0,05$.

Якщо емпіричне значення критерію виявляється на межі зони значущості, $p \leq 0,01$, у зоні значущості – що $p < 0,01$.

Підрахунок критерію Q Розенбаума

1) Перевірити, чи виконуються обмеження: $n_1, n_2 \geq 11$, $n_1 \approx n_2$.

2) Упорядкувати значення окремо в кожній вибірці за ступенем зростання ознаки. Вважати вибіркою 1 вибірку, значення в якій імовірно вищі, а вибіркою 2 – значення якої імовірно нижчі.

3) Визначити найвище (максимальне) значення у вибірці 2.

4) Підрахувати кількість значень у вибірці 1, які вищі за максимальне значення у вибірці 2. Позначити отриману величину як S_1 .

5) Визначити найнижче (мінімальне) значення у вибірці 1.

6) Підрахувати кількість значень у вибірці 2, які нижчі за мінімальне значення вибірки. Позначити отриману величину як S_2 .

7) Підрахувати емпіричне значення Q за формулою:
 $Q_{\text{емп.}} = S_1 + S_2$

8) За таблицею визначити критичні значення Q для даних n_1 і n_2 . Якщо $Q_{\text{емп.}}$

Дорівнює $Q_{0,05}$, або перевищує його, H_0 відкидається.

9) За $n_1, n_2 > 26$ зіставити отримане емпіричне значення з $Q_{\text{кр.}} = 8$ ($p \leq 0,05$) і $Q_{\text{кр.}} = 10$ ($p \leq 0,01$). Якщо $Q_{\text{емп.}}$ перевищує або принаймні дорівнює $Q_{\text{кр.}} = 8$, H_0 відкидається.

6.4 Т-критерій Вілкоксона

Призначення критерію.

Критерій застосовується для зіставлення показників, виміряних у *двох різних умовах на одній і тій самій вибірці випробовуваних*.

Він дає змогу встановити не тільки спрямованість змін, а й їхню вираженість. З його допомогою ми визначаємо, чи є зрушення показників у якомусь одному напрямку більш інтенсивнішим, ніж в іншому.

Опис критерію T.

Цей критерій застосовний у тих випадках, коли ознаки виміряні принаймні за шкалою порядку, і зсуви між другим і першим вимірами теж можуть бути впорядковані. Для цього вони мають варіювати в досить широкому діапазоні. У принципі, можна застосовувати критерій T і в тих випадках, коли зсуви приймають тільки три значення: -1, 0 і +1, але тоді критерій T навряд чи додасть що-небудь нове до тих висновків, які можна було б отримати за допомогою критерію знаків. Ось якщо зрушення змінюються, скажімо, від -30 до +45, тоді має сенс їх ранжувати і потім підсумовувати ранги.

Суть методу полягає в зіставленні вираженості зрушень у тому й іншому напрямках за абсолютною величиною. Для цього спочатку ранжуються всі абсолютні величини зрушень, а потім підсумовуються ранги. Якщо зрушення в позитивний і в негативний бік відбуваються випадково, то суми рангів абсолютних значень їх будуть приблизно рівні. Якщо ж

інтенсивність зсуву в одному з напрямків переважає, то сума рангів абсолютних значень зрушень у протилежний бік буде значно нижчою, ніж це могло б бути за випадкових змін.

Спочатку виходять із припущення про те, що типовим зсувом буде зсув у напрямку, що частіше зустрічається, а нетиповим, або рідкісним, зсувом – зсув у напрямку, що зустрічається більш рідко.

Гіпотези.

H_0 : Інтенсивність зсувів у типовому напрямку не перевищує інтенсивності зсувів у нетиповому напрямку.

H_1 : Інтенсивність зрушень у типовому напрямку перевищує інтенсивність зрушень у нетиповому напрямку.

Обмеження в застосуванні критерію T Вілкоксона

1) Мінімальна кількість випробовуваних, які пройшли вимірювання у двох умовах – 5 осіб. Максимальна кількість випробовуваних – 50 осіб, що диктується верхньою межею наявних таблиць. Критичні значення T наведено в таблиці.

2) Нульові зсуви з розгляду виключаються, і кількість спостережень n зменшується на кількість цих нульових зсувів (McCall R., 1970, р. 36). Можна обійти це обмеження, сформулювавши гіпотези, що включають відсутність змін, наприклад: «Зсув у бік збільшення значень перевищує зсув у бік зменшення значень і тенденцію збереження їх на колишньому рівні».

Приклад 6.5. У вибірці курсантів військового училища (юнаки віком від 18 до 20 років) вимірювали здатність до утримання фізичного вольового зусилля на динамометрі. Спочатку у випробовуваних вимірювали максимальну м'язову силу кожної з рук, а наступного дня їм пропонувалося витримувати на динамометрі з рухомою стрілкою м'язове зусилля, що дорівнює $1/2$ максимальної м'язової сили даної руки. Відчувши втому, випробуваний мав повідомити про це

експериментатора, але не припиняти дослід, долаючи втому і неприємні відчуття – «боротися, поки воля не вичерпається». Дослід проводили двічі; спочатку зі звичайною інструкцією, а потім, після того, як випробуваний заповнював опитувальник самооцінки вольових якостей за методикою А.Ц. Пуні, йому пропонувалося уявити собі, що він уже домогся ідеалу в розвитку вольових якостей, і продемонструвати відповідне ідеалу вольове зусилля. Чи підтвердилася гіпотеза експериментатора про те, що звернення до ідеалу сприяє зростанню вольового зусилля? Дані подано в таблиці 6.6.

Таблиця 6.6

Код імені випробуваного		Час утримання зусилля на динамометрі		Різниця ($t_{\text{після}}$ – $t_{\text{до}}$)	Абсолютне значення різниці	Ранговий номер різниці
		До вимір. вольових якостей ($t_{\text{до}}$)	Після вимір. вольових якостей ($t_{\text{після}}$)			
1	В.	64	25	-39	39	11
2	Кап.	77	50	-27	27	8
3	Кор.	69	72	+3	3	1
4	Кут.	95	76	-19	19	6
5	Л.	105	67	-38	38	9,5
6	М.	83	75	-8	8	4
7	Р.	73	77	+4	4	2,5
8	С.	75	71	-4	4	2,5
9	Т.	101	63	-38	38	9,5
10	Х.	97	122	+25	25	7
11	Ю	78	60	-18	18	5
						Сума: 66

Для підрахунку цього критерію немає необхідності впорядковувати ряди значень за наростанням ознаки. Можна використовувати алфавітний список випробовуваних, як у цьому випадку.

Перший крок у підрахунку критерію T – віднімання кожного індивідуального значення «до» зі значення «після» (можна віднімати значення «після» зі значень «до», це ніяк не вплине на розрахунок критерію. Але краще в усіх випадках дотримуватися однієї системи, щоб не заплутатися самим).

Ми бачимо з табл., що 8 отриманих різниць від'ємні і лише 3 позитивні.

Це означає, що у 8 випробовуваних тривалість утримання м'язового зусилля в другому вимірі зменшилася, а у 3 – збільшилася. Ми зіткнулися з тим випадком, коли вже зараз не можна сформулювати статистичну гіпотезу, яка відповідає первісному припущенню дослідника. Передбачалося, що звернення до ідеалу буде збільшувати тривалість м'язового зусилля, а експериментальні дані свідчать, що лише в 3 випадках з 11 цей показник дійсно збільшився. Можна сформулювати лише гіпотезу, яка передбачає не суттєвості зсуву цього показника в бік зниження.

Гіпотези.

H_0 : *Інтенсивність зрушень у бік зменшення тривалості м'язового зусилля не перевищує інтенсивності зрушень у бік її збільшення.*

H_1 : *Інтенсивність зрушень у бік зменшення тривалості м'язового зусилля перевищує інтенсивність зрушень у бік її збільшення.*

На наступному кроці всі зрушення, незалежно від їхнього знаку, мають бути проранжовані за вираженості. У таблиці у четвертому зліва стовпчику наведено абсолютні величини зсувів, а в останньому стовпчику (праворуч) – ранги цих абсолютних величин. Меншому значенню відповідає менший ранг. При цьому сума рангів дорівнює 66, що відповідає розрахунковій:

$$\sum R_i = \frac{N(N+1)}{2} = \frac{11 \cdot 12}{2} = 66.$$

Тепер відзначимо ті зрушення, які є нетиповими, у даному випадку – позитивними. У таблиці ці зрушення і відповідні їм ранги виділено кольором.

Сума рангів цих «рідкісних» зрушень становить емпіричне значення критерію T :

$$T_{\text{емп.}} = \sum R_r = 1 + 2,5 + 7 = 10,5$$

де R_r – рангові значення зсувів з більш нетиповим знаком.

За таблицею знаходимо критичні значення для T -критерію Вілкоксона для $n = 11$:

$$T_{\text{кр.}} = \begin{cases} 13, & (p \leq 0,05) \\ 7, & (p \leq 0,01) \end{cases} \Rightarrow$$

$$T_{\text{кр.}} = 7 (p \leq 0,01) < T_{\text{емп.}} = 10,5 < T_{\text{кр.}} = 13 (p \leq 0,05)$$

Зона значущості в цьому разі простягається вліво, справді, якби «рідкісних», у даному разі позитивних, зрушень не було зовсім, то й сума їхніх рангів дорівнювала б нулю. У цьому ж випадку емпіричне значення T потрапляє в зону невизначеності (рис. 6.2):

$$T_{\text{емп.}} = 10,5 < T_{\text{кр.}} = 13 (p \leq 0,05)$$

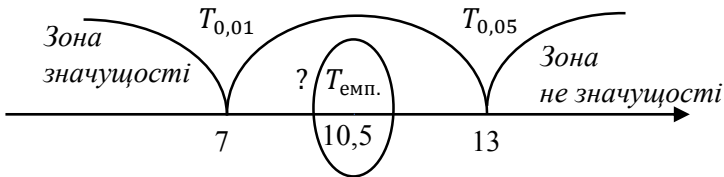


Рисунок 6.2

Відповідь: H_0 підтверджується. *Інтенсивність негативного зсуву показника фізичного вольового зусилля перевищує інтенсивність позитивного зсуву ($p \leq 0,05$)*

Таким чином, тривалість утримання м'язового вольового зусилля в другому вимірі знижується, і це зрушення не випадкове. Інструкція, що орієнтує випробовуваного на відповідність ідеалу в розвитку волі, виявилася набагато менш потужним чинником, ніж якась інша сила – можливо, м'язова втома, можливо, розчарування в собі або в можливостях цього психологічного експерименту. А можливо, в момент другого виміру просто перестає діяти якийсь потужний фактор, який був активний спочатку?

Підрахунку критерію T Вілкоксона

1) Скласти список випробовуваних у будь-якому порядку, наприклад, алфавітному.

2) Обчислити різницю між індивідуальними значеннями в другому і першому вимірах («після» – «до»). Визначити, що вважатиметься «типовим» зсувом і сформулювати відповідні гіпотези.

3) Перевести різниці в абсолютні величини і записати їх окремим стовпчиком (інакше важко відволіктися від знака різниці).

4) Проранжувати абсолютні величини різниць, нараховуючи меншому значенню менший ранг. Перевірити збіг отриманої суми рангів із розрахунковою.

5) Зазначити кружками або іншими знаками ранги, що відповідають зсувам у «нетиповому» напрямку.

6) Підрахувати суму цих рангів за формулою:

$$T_{\text{емп.}} = \sum R_r$$

де R_r – рангові значення зсувів з більш нетиповим знаком.

7) Визначити критичні значення T для даного n за таблицею критичних значень

Якщо $T_{\text{емп.}}$ Менше або дорівнює $T_{\text{кр.}}$, зсув у «типовий» бік за інтенсивністю достовірно переважає.

6.5 Виявлення відмінностей у розподілі ознаки.

χ^2 критерій Пірсона

Розподіли можуть відрізнятися за середніми, дисперсіями, асиметрією, ексцесом і за поєднанням цих параметрів. Розглянемо кілька прикладів.

На рисунку 6.3 подано два розподіли ознаки. Розподіл 1 характеризується меншим діапазоном варіативності та меншою дисперсією, ніж розподіл 2. В розподілі 1 частіше зустрічаються

значення ознаки, близькі до середньої, а в розподілі 2 частіше трапляються вищі та нижчі, ніж середні значення ознаки.

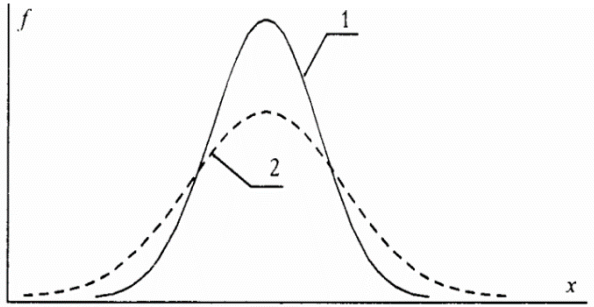


Рисунок 6.3

Саме таке співвідношення може спостерігатися в розподілі фенотипічних ознак у чоловіків (крива 2) і жінок (крива 1). Фенотипічна дисперсія чоловічої статі має бути більшою, ніж жіночої. Чоловіки – це авангардна частина популяції, відповідальна за пошук нових форм пристосування, тому в них частіше зустрічаються рідкісні крайні значення різних фенотипічних ознак. Ці відхилення, носять «футуристичний» характер, це «проби», що включають як майбутні можливі шляхи еволюції, так і помилки. Водночас жіноча частина популяції відповідальна за збереження вже накопичених змін, тому в них частіше зустрічаються середні значення фенотипічних ознак.

Аналіз реально одержуваних у дослідженнях розподілів може дозволити нам підтвердити або спростувати ці теоретичні припущення.

Важку задачу більшість випробовуваних розв’язують у тенденції довше, ніж просту, але в той же час майже завжди знаходяться люди, які вирішують його миттєво. Якщо ми доведемо, що розподіли статистично достовірно різняться, це може стати основою для побудови класифікацій завдань і типологій випробовуваних.

Часто буває корисно зіставити отриманий емпіричний розподіл із теоретичним розподілом. Наприклад, для того, щоб довести, що він підпорядковується або, навпаки, не підпорядковується нормальному закону розподілу.

У практичних цілях емпіричні розподіли повинні перевірятися на «нормальність» у тих випадках, коли ми маємо намір використовувати параметричні методи і критерії.

Традиційні для вітчизняної математичної статистики критерії визначення розбіжності або згоди розподілів – це метод χ^2 -критерій Пірсона та λ -критерій Колмогорова-Смирнова.

Обидва ці методи вимагають ретельного групування даних і досить складних обчислень. Крім того, можливості цих критеріїв повною мірою проявляються на великих вибірках ($n > 30$). Проте, вони можуть виявитися настільки незамінними, що досліднику доведеться знехтувати економією часу і зусиль. Наприклад, вони незамінні в таких двох випадках:

1) у задачах, що потребують доказу не випадковості переваг у виборі з кількох альтернатив;

2) у завданнях, що вимагають виявлення точки максимальної розбіжності між двома розподілами, яку потім використовують для перегрупування даних з метою застосування критерію φ^* (кутового перетворення Фішера).

Призначення критерію.

Критерій χ^2 застосовують із двома цілями:

1) для зіставлення емпіричного розподілу ознаки з теоретичним – рівномірним, нормальним або якимось іншим;

2) для зіставлення двох трьох або більше емпіричних розподілів однієї й тієї ж ознаки. Насправді сфери застосування критерію χ^2 різноманітні, але в цьому посібнику ми обмежуємося тільки цими двома, що найчастіше зустрічаються на практиці, цілями).

Опис критерію.

Критерій χ^2 відповідає на запитання про те, чи з однаковою частотою зустрічаються різні значення ознаки в

емпіричному і теоретичному розподілах або у двох і більше емпіричних розподілах.

Перевага методу полягає в тому, що він дає змогу зіставляти розподіли ознак, представлених у будь-якій шкалі, починаючи від шкали найменувань. У найпростішому випадку альтернативного розподілу «так – ні», «допустив брак – не допустив браку», «вирішив завдання – не вирішив завдання».

Під час зіставлення емпіричного розподілу з теоретичним ми визначаємо ступінь розбіжності між емпіричними і теоретичними частотами.

У разі зіставлення двох емпіричних розподілів ми визначаємо ступінь розбіжності між емпіричними частотами і теоретичними частотами, які спостерігалися б у разі збігу двох цих емпіричних розподілів. Формули розрахунку теоретичних частот будуть спеціально дані для кожного варіанта зіставлень.

Що більша розбіжність між двома зіставляваними розподілами, то більшим є емпіричне значення χ^2 .

Гіпотези.

Можливі кілька варіантів гіпотез, залежно від завдань, які ми перед собою ставимо.

Перший варіант:

H_0 : Отриманий емпіричний розподіл ознаки не відрізняється від теоретичного (наприклад, рівномірного) розподілу.

H_1 : Отриманий емпіричний розподіл ознаки відрізняється від теоретичного розподілу.

Другий варіант:

H_0 : Емпіричний розподіл 1 не відрізняється від емпіричного розподілу 2.

H_1 : Емпіричний розподіл 1 відрізняється від емпіричного розподілу 2.

Третій варіант:

H_0 : Емпіричні розподіли 1, 2, 3, не відрізняються між собою.

H_1 : Емпіричні розподіли 1, 2, 3, розрізняються між собою.

Критерій χ^2 дає змогу перевірити всі три варіанти гіпотез. Обмеження критерію.

1) Обсяг вибірки має бути досить великим: $n > 30$. За $n < 30$ критерій χ^2 дає досить наближені значення. Точність критерію підвищується за великих n .

2) Теоретична частота для кожного осередку таблиці не повинна бути меншою за 5: $f \geq 5$. Це означає, що якщо число розрядів задано заздалегідь і не може бути змінено, то ми не можемо застосовувати метод, χ^2 не накопичивши певного мінімального числа спостережень. Якщо, наприклад, ми хочемо перевірити наші припущення про те, що частота звернень до телефонну службу Довіри нерівномірно розподіляються по 7 днях тижня, то нам буде необхідно $5 * 7 = 35$ звернень. Таким чином, якщо кількість розрядів (k) задано заздалегідь, як у цьому випадку, мінімальне число спостережень (n_{min}) визначається за формулою: $n_{min} = k * 5$.

3) Обрані розряди мають «вичерпувати» весь розподіл, тобто охоплювати весь діапазон варіативності ознак. При цьому групування на розряди має бути однаковим у всіх розподілах, що зіставляються.

4) Необхідно вносити «поправку на безперервність» при зіставленні розподілів ознак, які приймають лише 2 значення. При внесенні поправки значення χ^2 зменшується (див. приклад із поправкою на безперервність).

5) Розряди мають бути такими, що не перехреснюються: якщо спостереження віднесено до одного розряду, то воно вже не може бути віднесене до жодного іншого розряду. Сума спостережень за розрядами завжди має дорівнювати загальній кількості спостережень.

Правомірним є питання про те, що вважати числом спостережень – кількість виборів, реакцій, дій або кількість випробовуваних, які роблять вибір, проявляють реакції або виконують дії. Якщо випробовуваний проявляє кілька реакцій, і всі вони реєструються, то кількість випробовуваних не буде рівним кількості реакцій. Ми можемо підсумувати реакції кожного випробовуваного, як, наприклад, це робиться в методиці Хекхаузена для дослідження мотивації досягнення або в тесті фрустраційної толерантності С. Розенцвейга, і порівнювати розподіли індивідуальних сум реакцій у кількох вибірках.

У цьому разі числом спостережень буде кількість випробовуваних. Якщо ж ми підраховуємо частоту реакцій певного типу в цілому по вибірці, то отримуємо розподіл реакцій різного типу, і в цьому випадку кількістю спостережень буде загальна кількість зареєстрованих реакцій, а не кількість випробовуваних.

З математичної точки зору правило незалежності розрядів дотримується в обох випадках: одне спостереження належить до одного і тільки одного розряду розподілу.

Але чи вважати спостереженням кожного випробовуваного або кожен досліджувану реакцію випробовуваного – питання, вирішення якого залежить від цілей дослідження.

Приклад 6.6 (з поправкою на неперервність).

У дослідженні порогів соціального атома професійних психологів просили визначити, з якою частотою зустрічаються в їхньому записнику чоловічі та жіночі імена колег-психологів. Спробуємо визначити, чи відрізняється розподіл, отриманий за записнику жінки-психолога Х, від рівномірного розподілу. Емпіричні частоти подано в таблиці 6.7

Таблиця 6.7

<i>Чоловіки</i>	<i>Жінки</i>	<i>Всього</i>
22	45	67

Гіпотези.

H_0 : Розподіл чоловічих і жіночих імен у записнику X не відрізняється від рівномірного розподілу.

H_1 : Розподіл чоловічих і жіночих імен у записній книжці X відрізняється від рівномірного розподілу.

Кількість спостережень $n = 67$; кількість значень ознаки $k = 2$. Розрахуємо *теоретичну частоту*:

$$f_{\text{теор.}} = n/k = 33,5$$

Число ступенів свободи $\nu = k - 1$.

Далі всі розрахунки проводимо за відомим алгоритмом, але з одним додаванням: перед зведенням у квадрат різниці частот ми повинні зменшити абсолютну величину цієї різниці на 0,5 (табл. 6.8, четвертий стовпець).

Таблиця 6.8.

<i>Розряди – приналежність до тієї чи іншої статі</i>	<i>Емпірична част. $f_{\text{емп.}}$</i>	<i>Теоретична част. $f_{\text{теор.}}$</i>	$f_{\text{емп.}} - f_{\text{т.}}$	$f_{\text{емп.}} - f_{\text{т.}} - 0,5$	$(f_{\text{емп.}} - f_{\text{т.}} - 0,5)^2$	$\frac{(f_{\text{емп.}} - f_{\text{т.}} - 0,5)^2}{f_{\text{т.}}}$
<i>Чоловіки</i>	22	33,5	-11,5	-11	121	3,61
<i>Жінки</i>	45	33,5	+11,5	+11	121	3,61

За таблицею критичних значень для $\nu = 1$ маємо:

$$\chi_{\text{кр}}^2 = \begin{cases} 3,841 & (p \leq 0,05) \\ 6,635 & (p \leq 0,01) \end{cases}$$
$$\chi_{\text{емп}}^2 = 7,22$$
$$\chi_{\text{емп}}^2 > \chi_{\text{кр}}^2$$

Відповідь: H_0 відхиляється і приймається H_1 . Розподіл чоловічих і жіночих імен у записній книжці Х відрізняється від рівномірного розподілу ($p \leq 0,01$).

Приклад 6.7 (порівняння двох емпіричних розподілів).
Визначимо, чи відрізняються розподіли чоловічих і жіночих імен у психолога А і психолога В, теж жінки. Емпіричні частоти наведено в таблиці 6.9

Таблиця 6.9

	Чоловіки	Жінки	Всього
Психолог А.	22 А.	45 Б.	67
Психолог В.	59 В.	109 Г.	168
Всього	81	154	235

Гіпотези.

H_0 : Розподіли чоловічих і жіночих імен у двох записниках не відрізняються.

H_1 : Розподіли чоловічих і жіночих імен у двох записниках різняться між собою.

Теоретичні частоти розраховуємо за вже відомою формулою:

$$t_{\text{теор.}} = \frac{\left(\begin{array}{c} \text{сума частот за} \\ \text{відповідним рядком} \end{array} \right) \cdot \left(\begin{array}{c} \text{сума частот за} \\ \text{відповідним стовпцем} \end{array} \right)}{\text{загальна кількість спостережень}}$$

$$f_{\text{А.теор.}} \cdot \frac{67 \cdot 81}{235} = 23,09;$$

$$f_{\text{Б.теор.}} \cdot \frac{67 \cdot 154}{235} = 43,91;$$

$$f_{\text{В.теор.}} \cdot \frac{168 \cdot 81}{235} = 57,91;$$

$$f_{\text{Г.теор.}} \cdot \frac{168 \cdot 154}{235} = 110,09.$$

Таблиця 6.10

Ячейки емпіричних частот	Емпірична частота $f_{\text{емп.}}$	Теоретична частота $f_{\text{теор.}}$	$f_{\text{емп.}}$ – $f_{\text{т.}}$	$f_{\text{емп.}}$ – $f_{\text{т.}}$ – 0,5	$(f_{\text{емп.}}$ – $f_{\text{т.}}$ – 0,5) ²	$\frac{(f_{\text{емп.}} - f_{\text{т.}} - 0,5)^2}{f_{\text{т.}}}$
А.	22	23,09	-1,09	0,59	0,35	0,015
Б.	45	43,91	+1,09	0,59	0,35	0,008
В.	59	57,91	+1,09	0,59	0,35	0,006
Г.	109	110,09	-1,09	0,59	0,35	0,003
	235	235,0	0			$\chi_{\text{емп}}^2 = 0,032$

За таблицею критичних значень для $\nu = 1$ маємо:

$$\chi_{\text{кр}}^2 = \begin{cases} 3,841 & (p \leq 0,05) \\ 6,635 & (p \leq 0,01) \end{cases}$$

$$\chi_{\text{емп}}^2 = 0,032$$

$$\chi_{\text{емп}}^2 < \chi_{\text{кр}}^2$$

Відповідь: Приймається гіпотеза H_0

Зауваження. Поправки на безперервність можна уникнути, якщо подібного роду завдання розв'язувати за допомогою ϕ^* -критерію Фішера.

Розрахунок критерію χ^2

1) Занести в таблицю найменування розрядів і відповідні їм емпіричні частоти (перший стовпець).

2) Поруч із кожною емпіричною частотою записати теоретичну частоту (другий стовпчик).

3) Підрахувати різниці між емпіричною та теоретичною частотою за кожним розрядом (рядку) і записати їх у третій стовпчик.

4) Визначити число ступенів свободи за формулою: $\nu = k - 1$, де k – кількість розрядів ознаки. Якщо $\nu = 1$, внести поправку на «безперервність».

5) Підвести до квадрата отримані різниці та занести їх у четвертий стовпчик.

6) Розділити отримані квадрати різниць на теоретичну частоту і записати результати в п'ятий стовпець.

7) Просумувати значення п'ятого стовпчика. Отриману суму позначити як $\chi^2_{\text{емр}}$

8) Визначити за таблицею критичні значення для даного числа ступенів свободи ν . Якщо $\chi^2_{\text{емр}}$ менше критичного значення, розбіжності між розподілами статистично недостовірні.

Якщо $\chi^2_{\text{емр}}$ дорівнює критичному значенню або перевищує його, розбіжності між розподілами статистично достовірні.

ЛЕКЦІЯ 7

РЕГРЕСІЙНИЙ АНАЛІЗ

1. Задачі регресійного аналізу.
2. Визначення коефіцієнтів регресії.
3. Обчислення похибки рівняння регресії.
4. Види рівняння регресії.

7.1 Задачі регресійного аналізу

Раніше ми зазначали, що пошук зв'язку між явищами є однією з основних задач будь-якої науки. Однак, аналіз зв'язку має дві сторони:

1. Сила зв'язку і його напрям (прямий чи обернений).
2. Форма зв'язку (пряма, параболічна тощо).

Для аналізу сили зв'язку та його напрямку ми з Вами використовували кореляційний аналіз, і наголошували на тому, що досить часто зв'язок може носити досить складний нелінійний характер, а тому одного лише кореляційного аналізу недостатньо. Необхідно аналізувати діаграму розсіювання, на якій наочно видно, якої форми може бути зв'язок.

Але одного лише наочного зображення для переконливих висновків недостатньо. Статистика дозволяє нам з великою точністю визначити форму зв'язку та записати її у вигляді математичного рівняння з залежною (результативною) та незалежною (факторною) змінною. Методом, з допомогою якого можна це зробити є регресійний аналіз.

ЗАДАЧІ	МЕТОДИ ВИРІШЕННЯ
Пошук сили зв'язку та його напрямку	Кореляційний аналіз
Пошук форми зв'язку, передбачення змін	Регресійний аналіз

Найпростішим випадком є лінійна форма зв'язку, а тому розгляд регресійного аналізу ми почнемо саме з нього.

7.2 Визначення коефіцієнтів регресії

В одному з попередніх питань ми розглядали завдання – визначити зв'язок між кількістю учбовою успішністю учнів та їх особистісною тривожністю. Для вирішення цього завдання ми користувалися коефіцієнтом лінійної кореляції Пірсона (табл. 7.1), і отримали такий результат – $r_{xy} = 0,87$.

Таблиця 7.1

№	Учні	Успішність учнів (x)	Особистісна тривожність (y)	x_i^2	y_i^2	$x_i y_i$
1.	А.О.	10	48	100	2304	480
2.	Л.А.	12	49	144	2401	588
3.	М.П.	10	45	100	2025	450
4.	С.А.	9	38	81	1444	342
5.	Р.К.	8	34	64	1156	272
6.	В.О.	7	25	49	625	175
7.	Д.К.	10	42	100	1764	420
8.	В.К.	9	41	81	1681	369
9.	М.А.	11	43	121	1849	473
10.	Д.Н.	10	46	100	2116	460
11.	З.Ц.	12	48	144	2304	576
12.	В.Ж.	8	23	64	529	184
		$\sum x_i = 116$	$\sum y_i = 482$	$\sum x_i^2 = 1148$	$\sum y_i^2 = 20198$	$\sum x_i y_i = 4789$

Давайте спробуємо застосувати до цих же даних регресійний аналіз і встановимо форму та рівняння зв'язку між тривожністю та успішністю. Дійсно, результати кореляційного аналізу дають нам можливість сказати лише, що чим вища успішність, тим вища тривожність (або навпаки, нижча тривожність, тим нижча успішність), але ми не знаємо, на скільки зростає одна змінна при зменшенні іншої – на скільки зростає тривожність при зростанні успішності на 1 бал? Або: на скільки зростає успішність при зростанні тривожності на 1 бал?

Приклад 7.1 Перш за все, необхідно побудувати діаграму розсіювання, використовуючи дані таблиці 7.1. Нагадаємо, що у регресійному аналізі, як і в кореляційному, позначання однією змінної x , а другої y – це умовність. Що на що впливатиме ми не будемо знати до тих пір, поки не проведемо експеримент!!!

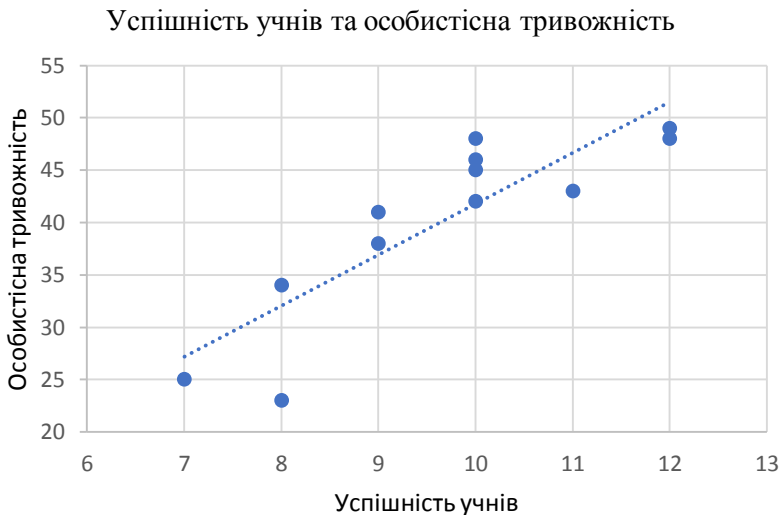


Рисунок 7.1

Нашою задачею буде знайти лінію, яка максимально наблизатиметься до всіх точок діаграми розсіювання, і відповідно, наближено описуватиме всіх їх.

Загальне рівняння прямої лінії має вигляд:

$$Y = a + bX ,$$

де X – умовно незалежна змінна (факторна змінна), Y – умовно залежна змінна (результативна змінна) (умовно, бо ми на основі кореляційного та регресійного аналізу не можемо визначити, що є причиною, а що – наслідком).

Коефіцієнти **a** та **b** називають *коефіцієнтами регресії*, а саму лінію – лінією регресії.

7.2.1 Метод вибраних точок

У нашому випадку X – успішність учнів, Y – особистісна тривожність. Найпростішим, але найнедосконалішим методом регресійного аналізу є **метод вибраних точок**.

Суть методу вибраних точок полягає у тому, що на графіку обираються дві точки, так, щоб проведена через них лінія ділила все поле точок на приблизно однакові частини. У нашому випадку – це точки $(7; 25)$ та $(10; 46)$ – рисунок 7.1. Підставляючи значення цих точок у рівняння прямої отримують систему з двох рівнянь, з яких знаходять коефіцієнти загального рівняння a та b .

У нашому випадку:

$$\text{В. О. : } X = 7, Y = 25;$$

$$\text{Д. К. : } X = 10, Y = 46.$$

Відповідно, утворюється система рівнянь:

$$\begin{cases} 25 = a + 7b \\ 46 = a + 10b \end{cases}$$

З першого рівняння системи віднімемо друге:

$$-21 = -3b; \quad b = 7.$$

Підставимо отримане значення в одне з рівнянь, отримуємо $a = -24$.

А тому, загальне рівняння регресії для наших даних виглядатиме так:

$$Y = -24 + 7X$$

Маючи це рівняння, ми можемо вийти за межі отриманих нами конкретних даних, зробивши прогноз і з'ясувавши, як буде змінюватися тривожність учнів з підвищенням чи зниженням їх успішності.

Перевагою методу вибраних точок є його швидкість, а недоліком – неточність та приблизність. Він не враховує всіх значень, отриманих у дослідженні, а тому його використання обмежене первинним аналізом.

7.2.2 Метод найменших квадратів

Розібравшись у найпростішому методі, можемо перейти до складнішого, але точнішого – **методу найменших квадратів**.

Проведемо по нашій діаграмі розсіювання нову уявну лінію регресії (уявну, бо саме рівняння цієї лінії ми будемо шукати, це ніби ескіз). Для кожної реальної експериментальної точки можна на лінії побудувати теоретичну точку. Відповідно, рівняння цієї теоретичної кривої матиме вигляд:

$$Y_{\text{теор}} = a + bX,$$

Де $Y_{\text{теор}}$ – теоретична ордината.

Відповідно, можна знайти різницю між теоретичними та експериментальними ординатами кожної точки:

$$\begin{cases} Y_{\text{теор } 1} - Y_1 = a + bX_1 - Y_1 \\ \dots \\ Y_{\text{теор } 10} - Y_{10} = a + bX_{10} - Y_{10} \end{cases}.$$

Розв'язати таку систему рівнянь можна, однак, у процесі її розв'язку можна стикнутися з такою складністю – внаслідок того, що різниця $Y_{\text{теор } 1} - Y_1$ може бути як додатною, так і від'ємною їх сума може бути рівною нулю. Тому вирішили знаходити не суми різниць експериментальних і теоретичних ординат, а суми квадратів різниць цих ординат – так зникає вплив знаків. Отже, вирішення поставленої задачі зводиться до отримання такої лінії, для якої сума квадратів відхилень експериментальних значень від обчислених теоретично є найменшою. У результаті розв'язку цих рівнянь (методами диференціального числення) отримується така загальна система:

$$\begin{cases} aN + b \sum X = \sum Y \\ a \sum X + b \sum X^2 = \sum XY \end{cases}$$

Отже, для пошуку лінії регресії методом найменших квадратів слід знайти всі необхідні суми, підставити їх у систему

рівнянь і розв'язати цю систему – коефіцієнти будуть знайдені. Для нашого випадку це буде таке рівняння:

$$\begin{cases} 12a + 116b = 482 \\ 116a + 1148b = 4789 \end{cases}$$

Розв'язавши його, ми отримаємо: $a = -39,49$; $b = 8,24$. А тому, маємо загальне рівняння регресії:

$$Y = -39,49 + 8,24X$$

Як бачимо, між методом вибраних точок та методом найменших квадратів існують певні відмінності.

7.3 Обчислення похибки рівняння регресії

Рівняння регресії дає можливість знайти теоретичне значення Y , яке іноді називають обчисленим значенням, вирівняним значенням або найбільш ймовірним значенням.

Лінія регресії не точно описує всю сукупність даних – за означенням вона максимально наближається до всіх точок діаграми розсіювання. Тому існує можливість відхилення даних від лінії регресії. Чим більше це відхилення, тим менш ефективною є регресія, і навпаки, чим менше це відхилення, тим регресія більш очевидна і ефективна. Величину, яка характеризує похибку рівняння регресії (середню квадратичну похибку рівняння, вибіркове стандартне відхилення від лінії регресії) шукають так:

1. Обчислити всі теоретичні значення "залежної" змінної за формулою регресії.
2. Знайти різниці між теоретичними та експериментальними значеннями "залежної" змінної.
3. Піднести різниці до квадрату та знайти їх суму.
4. Поділити суму квадратів на $n - 2$ (де n – об'єм вибірки).
5. Знайти квадратний корінь з отриманого числа.

Формула для обчислення похибки рівняння регресії виглядає так:

$$S_{yx} = \sqrt{\frac{\sum (Y_{\text{теор } i} - Y_i)^2}{n - 2}}.$$

Продовження прикладу 7.1

№	Учні	Усп-сть учнів (x)	Особ. Три-сть (y)	$Y_{\text{теор } i}$	$Y_{\text{теор } i} - Y_i$	$(Y_{\text{теор } i} - Y_i)^2$
1.	А.О.	10	48	42,91	-5,09	25,91
2.	Л.А.	12	49	59,39	10,39	107,95
3.	М.П.	10	45	42,91	-2,09	4,37
4.	С.А.	9	38	34,67	-3,33	11,09
5.	Р.К.	8	34	26,43	-7,57	57,30
6.	В.О.	7	25	18,19	-6,81	46,38
7.	Д.К.	10	42	42,91	0,91	0,83
8.	В.К.	9	41	34,67	-6,33	40,07
9.	М.А.	11	43	51,15	8,15	66,42
10.	Д.Н.	10	46	42,91	-3,09	9,55
11.	З.Ц.	12	48	59,39	11,39	129,73
12.	В.Ж.	8	23	26,43	3,43	11,76

$$\sum (Y_{\text{теор } i} - Y_i)^2 = 511,36;$$

$$S_{yx} = \sqrt{\frac{\sum (Y_{\text{теор } i} - Y_i)^2}{n - 2}} = \sqrt{\frac{511,36}{10}} = 7,15.$$

Можна знайти показник ефективності регресії, поділивши стандартне відхилення по Y (S_y) на похибку регресії (S_{yx}): $8,73/7,15 = 1,22$ або 122%.

Як і у випадку з дисперсією та стандартним відхиленням, значення похибки регресії саме по собі мало що каже. Воно проявляється лише у порівнянні. Зокрема, можна підрахувати похибку регресії для методу вибраних точок і порівняти з похибкою, отриманою при аналізі методом найменших квадратів.

Те рівняння, яке буде мати меншу похибку, і буде більш точно відображати дійсність.

7.4 Види рівняння регресії

Ми зазначили, що лінія регресії може бути не лише прямою, але і кривою. Використовують ряд типових кривих регресії, для яких розроблені свої рівняння.

1. Рівняння гіперболічного виду:

$$Y = a + \frac{b}{X}$$

Його коефіцієнти можна визначити з такої системи:

$$\begin{cases} \sum YX = a \sum X + bN \\ \sum YX^2 = a \sum X^2 + b \sum X \end{cases}$$

2. Рівняння логарифмічної параболи: $Y = a + b \cdot \lg X$

Його коефіцієнти можна визначити з такої системи:

$$\begin{cases} \sum Y = aN + b \sum \lg X \\ \sum Y \cdot \lg X = a \sum \lg X + b \sum \lg^2 X \end{cases}$$

3. Рівняння степеневої функції: $Y = aX^b$

Його коефіцієнти можна визначити з такої системи:

$$\begin{cases} \sum \lg Y = N \lg a + b \sum \lg X \\ \sum \lg Y \cdot \lg X = \lg a \sum \lg X + b \sum \lg^2 X \end{cases}$$

Загальні принципи вибору рівняння регресії

Вибір певної форми регресії на різних етапах дослідження детермінується різними чинниками.

ЕТАП ДОСЛІДЖЕННЯ	СПОСІБ ВИБОРУ РІВНЯННЯ РЕГРЕСІЇ
Попередній аналіз	На основі досвіду та інтуїції
Поглиблений аналіз	1. Будуються декілька ліній регресії. 2. Обчислюються похибки регресій та перевіряються значимості коефіцієнтів регресії. 3. Обирається лінія регресії з мінімальною похибкою та максимальною значимістю коефіцієнтів.

Хоча на перший погляд регресійний аналіз виглядає досить громіздким і довгим, сьогодні, з використанням комп'ютерної техніки, він проводиться буквально за секунди. Однак, ручний варіант обробки даних є просто необхідним для розуміння суті регресійного аналізу, знання його особливостей, переваг та недоліків.

ЛЕКЦІЯ 8

ДИСПЕРСІЙНИЙ АНАЛІЗ

1. Поняття про факторні гіпотези.
2. Типи факторних планів.
3. Результати факторних експериментів.
4. Двофакторний дисперсійний аналіз для незалежних вибірок.
5. Проведення двофакторного дисперсійного аналізу. Критерій Фішера F .

8.1 Поняття про факторні гіпотези

Інколи достатньо складно, а то й неможливо сформулювати гіпотезу про вплив одного явища на інше. Причина цього – системна та комплексна взаємодія різних чинників, які накладаються один на другого, утворюють складну і багатогранну картину оточуючої нас реальності. Тому простих гіпотез про вплив одного явища на інше буває мало для розуміння цієї реальності – необхідні складніші підходи.

Одним із таких підходів є факторний експеримент. Для статистичної обробки результатів факторного експерименту і використовують дисперсійний аналіз.

Отже, факторний експеримент використовують у випадках, коли необхідно перевірити складні гіпотези про зв'язок між змінними. Загальний вигляд складної гіпотези такий:

Якщо $A_1, A_2, \dots, A_n$ то B

При цьому, $A_1, A_2, \dots, A_n$ – незалежні змінні, а B – залежна змінна. Якщо користуватися термінологією факторних експериментів, то незалежні змінні варто називати факторними ознаками, або просто факторами, а залежні – результуючими ознаками.

Існують декілька типів факторних гіпотез, які перевіряються з допомогою дисперсійного аналізу:

перший тип – гіпотези про окремий вплив кожного із факторів,

другий тип – гіпотези про взаємодію факторів (коли присутність одного з факторів змінює вплив іншого).

При цьому, відношення між факторами можуть бути різноманітні:

- відношення диз'юнкції;
- відношення кон'юнкції;
- лінійна залежність;
- адитивність;
- мультиплікативність тощо.

8.2 Типи факторних планів

Суть факторного планування експерименту полягає в тому, що усі рівні незалежних змінних (факторів) необхідно комбінувати один з одним. При цьому необхідно, щоб кількість експериментальних груп була рівна кількості комбінацій рівнів незалежних змінних.

Першою основою для виділення типів факторних планів є кількість незалежних змінних (факторів). Відповідно, виділяють такі типи планів:

- план для **двох факторів**;
- план для **трьох факторів**.

Використання одночасно більше трьох факторів для планування експерименту є недоцільним, оскільки навіть обробка результатів трьохфакторного експерименту є достатньо громіздкою.

Другою основою для типології є кількість рівнів кожного з факторів:

- план для двох рівнів фактора;
- план для трьох рівнів фактора.

Узагальнюючи наведені типології, можна побудувати таблицю 8.1, в якій подані типи факторних планів для двох факторів.

Таблиця 8.1

		Фактор 1	
		2 рівні	3 рівні
Фактор 2	2 рівні	2x2	3x2
	3 рівні	2x3	3x3

Найтиповішими є плани типу 2x2 та 3x3. Факторний план 2x2 дозволяє виявити ефект впливу двох факторів на результуючу ознаку, а план 3x3, крім того, дає можливість визначити форму зв'язку між змінними.

Очевидно, що два рівні фактора можна позначити, наприклад, як “низький” та “високий”, або як “є ефект” та “немає ефекту”. План 2x2 має такий вигляд (табл. 8.2)

Таблиця 8.2

		Фактор 1	
		Є вплив	Немає вплив
Фактор 2	Є вплив	A_1	A_2
	Немає впливу	A_3	A_4

У таблиці A_1, A_2, A_3, A_4 – досліджувані групи, кожна з яких потрапляє в свої умови впливу фактора. Група A_1 перебуває в умовах впливу обох факторів, група A_2 – в умовах впливу лише другого фактора, група A_3 – в умовах впливу лише першого фактора, група A_4 – в умовах відсутності впливу обох факторів.

Факторний план 3x3 також достатньо часто використовують, особливо тоді, коли фактори мають складнішу градацію, ніж “є вплив – немає впливу”. В цьому випадку факторний план має такий вигляд, який представлено в таблиці 8.3.

Відповідно, тут також $A_1, A_2, \dots, A_9$ – досліджувані групи, кожна з яких перебуває у своїх експериментальних умовах.

Для трьох факторів можливі різноманітні варіанти, типу 2х3х2, 2х3х3, 3х2х3 тощо.

Таблиця 8.3

		Фактор 1		
		Є вплив	Середній рівень впливу	Немає вплив
Фактор 2	Є вплив	A_1	A_2	A_3
	Середній рівень впливу	A_4	A_5	A_6
	Немає впливу	A_7	A_8	A_9

Як бачимо, із зростанням кількості факторів або рівнів факторів хоча б на один, план експерименту значно ускладняється. Дійсно, навіть план 3х3 вже достатньо складно реалізувати – необхідно 9 незалежних вибірок досліджуваних.

У зв'язку з цим можна виділити ще одну, **третю, основу для класифікації** типів факторних планів – спосіб підбору досліджуваних груп:

- **збалансований** план;
- **ротаційний** план.

У **збалансованих** планах кожній комбінації рівнів факторів піддається окрема досліджувана група. недолік цього плану – велика кількість груп, і відповідно, ускладнення процедури експерименту.

У **ротаційних** планах цей недолік знімається – тут всі досліджувані піддаються впливу всіх комбінацій рівнів факторів. Однак, на зміну одному недоліку приходить інший – з'являється ефект тренування. Дійсно, при такій побудові експерименту кожен попередній експериментальний вплив залишає свій слід у психіці досліджуваного, відповідно, впливаючи на його поведінку в нових умовах. Накопичуючись, ці впливи можуть

повністю спотворити результати дослідження. Про способи контролю цих сторонніх чинників можна довідатися з рекомендованої літератури.

Розглянемо, у якій формі постають перед дослідником результати факторних експериментів, та як можна їх інтерпретувати.

8.3 Результати факторних експериментів

Перше, що повинен зробити дослідник після проведення факторного експерименту, це представити його результати графічно. Це дасть змогу уявити тип взаємодії між факторами, а також спланувати напрямок подальшого дисперсійного аналізу та його доцільність. Дійсно, якщо графічне зображення говорить про відсутність будь-яких ефектів, навряд чи варто проводити подальший аналіз. З іншого боку, видимий графічно ефект насправді може бути лише видимістю, – тому для підтвердження дисперсійний аналіз просто необхідний.

Виділяють такі види взаємодії:

- 1) нульова;
- 2) перехресна;
- 3) розбіжна.

Для ілюстрації видів взаємодії наведемо приклад дослідження.

Приклад 8.1. Для дослідження впливу рівня інтелектуального розвитку та стану здоров'я на статус дитини у групі було проведено дослідження за факторним планом. Визначався рівень інтелекту, загальний стан здоров'я та соціометричний статус дітей.

У результаті аналізу результатів такого дослідження можна отримати різні форми взаємодії факторів.

Нульова взаємодія має такий вигляд (рис. 8.1)

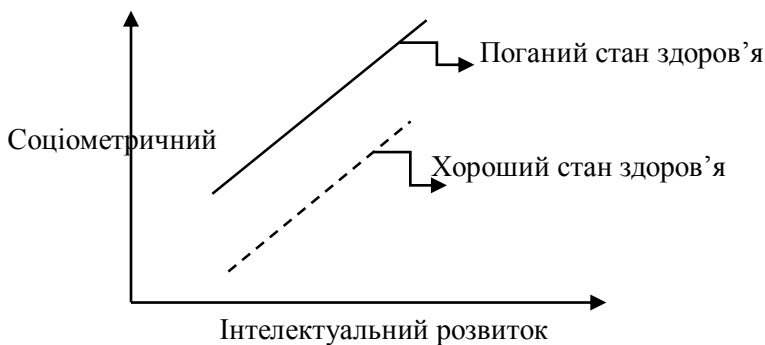


Рисунок 8.1

При нульовій взаємодії між факторами величина впливу другого фактора на результуючу ознаку однакова при всіх рівнях першого фактора – між факторами взаємодії немає.

Перехресна взаємодія має такий вигляд (рис. 8.2.):

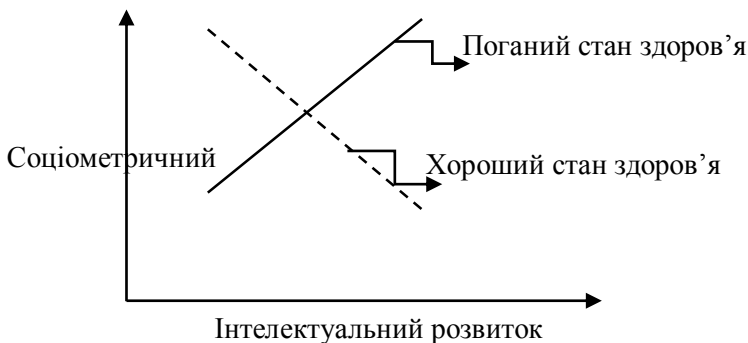


Рисунок 8.2

При перехресній взаємодії є можливість наочно спостерігати, в яких взаємозв'язках знаходяться фактори між собою. У зображеному випадку виявляється, що статус дитини в групі підвищується при низькому інтелекті та хорошому здоров'ї, а також при високому рівні інтелекту та поганому здоров'ї.

Розбіжна взаємодія має такий вигляд (рис. 8.3.):

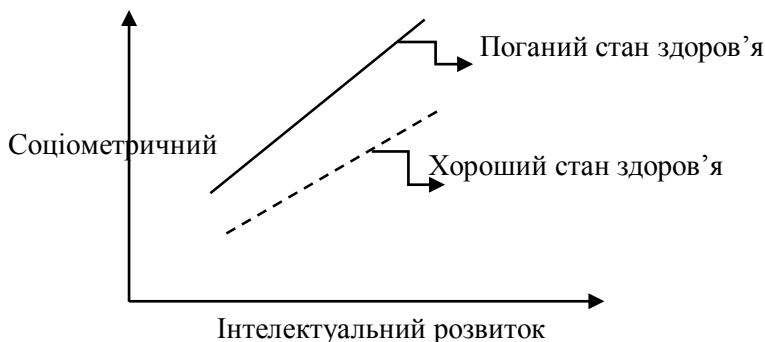


Рисунок 8.3

У випадку зображеної розбіжної взаємодії фактор «здоров'я» підсилює дію фактора «інтелект» — при наявності високого інтелекту погане здоров'я підвищує статус дитини. Існують також інші варіанти взаємодії між факторами, однак, у всіх випадках слід пам'ятати одне — якою б наочною не була графічна картина, проведення дисперсійного аналізу є обов'язковим.

8.4 Двофакторний дисперсійний аналіз для незалежних вибірок

1. Задача двофакторного дисперсійного аналізу

При аналізі результатів двофакторного експерименту варто послуговуватися методом двофакторного дисперсійного аналізу. Він дозволяє не лише оцінити вплив кожного із факторів на результуючу ознаку, але й оцінити їх взаємодію. Необхідність врахування взаємодії факторів можна побачити на таких прикладах.

Приклад 8.2. Дослідження впливу часового обмеження на швидкість вирішення задач може не вдатися без врахування фактору інтелекту. Може, наприклад, виявитися, що при

скороченні часу досліджувані з високим рівнем інтелекту будуть покращувати результати, а досліджувані з низьким – навпаки, будуть їх знижувати.

Приклад 8.3. В умовах жорсткого контролю за результатами діяльності внутрішньо мотивовані працівники можуть знижувати ефективність своєї діяльності, а мотивовані зовнішньо – підвищувати її.

Приклад 8.4. Залежність між мотивацією досягнення успіху та рівнем тривожності може бути прямою для жінок, та оберненою для чоловіків.

2. Обмеження двофакторного дисперсійного аналізу

Для того, щоб дисперсійний аналіз було проведено правильно, необхідно враховувати певні обмеження та рекомендації:

- Кожен фактор повинен мати не менше двох градацій.
- В кожній комірці дисперсійного комплексу повинно бути не менше двох спостережень (про дисперсійний комплекс див. далі).
- Дисперсійний комплекс повинен бути врівноваженим.
- Дисперсійний комплекс повинен бути симетричним (кожній градації першого фактора повинна відповідати однакова кількість градацій другого фактора).
- Фактори повинні бути незалежними (між ними не повинно існувати значимого кореляційного зв'язку).
- Результуюча ознака повинна бути розподілена нормально.

3. Підготовка даних до дисперсійного аналізу

Підготовка даних до дисперсійного аналізу полягає у перевірці всіх вище наведених обмежень. Більшість із них можна перевірити без складних математичних обчислень, крім 5 та 6 застереження.

До початку експерименту з допомогою кореляційного дослідження знайти рівень зв'язку між незалежними факторами. Якщо виявиться, що зв'язок не значний, можна приступати до

проведення факторного експерименту. Інакше – результати експерименту можуть бути абсолютно неочікуваними.

Стосовно 6 застереження, то варто відмітити, що дисперсійний аналіз є методом параметричним, а тому перевірка нормальності розподілу результуючої ознаки є обов'язковою (див. Лекція 3).

8.5 Проведення двофакторного дисперсійного аналізу.

Критерій Фішера F

Якщо виявиться, що результуюча ознака розподілена нормально, можна розпочинати аналіз. Прикладом факторного плану типу 2x2 може бути дослідження, яке описують Р. Солсо, Х. Джонсон, К. Біл (Солсо Р., Джонсон Х., Біл М., 2001, с. 89).

*Тривалий час існувало дві гіпотези відносно поведінки у ситуації змушеного захисту поглядів, протилежних власним. **Перша** – гіпотеза підкріплення. Відповідно до неї, чим більшою буде винагорода за захист поглядів, тим з більшою готовністю досліджуваній змінить свої власні погляди. **Друга** – гіпотеза дисонансу. Відповідно до теорії когнітивного дисонансу, навпаки, чим меншою буде винагорода, тим швидше досліджуваній змінить свої погляди. Для перевірки обох цих гіпотез було побудовано модель експерименту 2x2. Студентам коледжу запропонували написати твір, спрямований на підтримку закону про обмеження студентського самоврядування (це суперечило їх власним поглядам). Було обрано 4 групи досліджуваних. Першій та другій групам просто наказували написати твір на цю тему, а третій та четвертій запропонували на вибір – писати чи не писати. Крім того, першій та третій групі заплатили по \$0,50, а другій та четвертій – по \$2,50. Таким чином, було два фактори – наявність ситуації вибору (є та немає) та рівень винагороди (низький та високий). Результуючою ознакою був ступінь позитивного ставлення до закону.*

Спробуємо змоделювати процедуру цього експерименту (табл. 8.4.)

Таблиця 8.4

		Фактор А: ВІНАГОРОДА	
		Мала	Велика
Фактор В: ВІБІР	Немає вибору	Група 1	Група 2
	Є вибір	Група 3	Група 4

Нехай отримані такі первинні результати по кожному з досліджуваних (всього 16 досліджуваних, по 4 на кожен експериментальну групу) (табл. 8.5):

Таблиця 8.5

		Фактор А: ВІНАГОРОДА	
		Мала	Велика
Фактор В: ВІБІР	Немає вибору	10, 9, 5, 6	6, 4, 2, 3
	Є вибір	5, 4, 2, 4	8, 6, 5, 6

Можна висунути такі гіпотезу: між факторами “ВІБІР” та “ВІНАГОРОДА” буде спостерігатися взаємодія: при великій винагороді ставлення стає позитивним тоді, коли відсутній вибір, а при малій винагороді ставлення стає позитивним тоді, коли є можливість вибору.

Для подальшого аналізу важливо дотримуватися чіткого алгоритму.

1 крок: побудова дисперсійного комплексу.

Дисперсійний комплекс має такий вигляд (табл. 8.6):

Таблиця 8.6

		Фактор А: ВІНАГОРОДА		T_B
		Мала	Велика	
Фактор В: ВІБІР	Немає вибору	30	15	$T_{B1} = 45$
	Є вибір	15	25	$T_{B2} = 40$
T_A		$T_{A1} = 45$	$T_{A2} = 40$	$T = 85$

Тут T_A та T_B – відповідні суми показників зміни поглядів.

2 крок: перевірка названих вище обмежень дисперсійного аналізу. Вимога симетричності дисперсійного комплексу полягає в тому, щоб у кожній його комірці була однакова кількість спостережень. У нашому випадку на кожну комірку припадає по 4 спостереження.

3 крок: графічне зображення результатів експерименту. Можна зображати сумарні величини по кожній з комірок, а можна – середні арифметичні (рис. 8.4.).

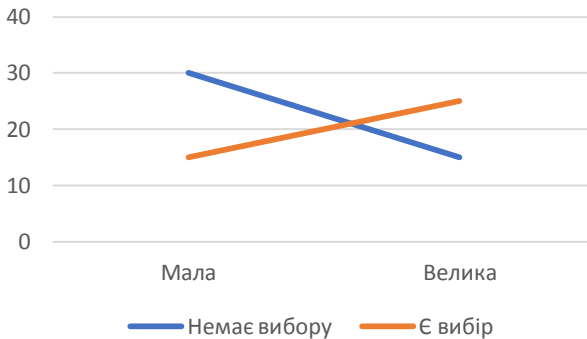


Рисунок 8.4

4 крок: побудова комплекту статистичних гіпотез.

Комплект №1: гіпотези про вплив фактора А

H_0 Відмінності у зміні ставлення до закону, обумовлені фактором А, є не більше вираженими, ніж випадкові відмінності.

H_1 Відмінності у зміні ставлення до закону, обумовлені фактором А, є більше вираженими, ніж випадкові відмінності.

Комплект №2: гіпотези про вплив фактора В

H_0 Відмінності у зміні ставлення до закону, обумовлені фактором В, є не більше вираженими, ніж випадкові відмінності.

H_1 Відмінності у зміні ставлення до закону, обумовлені фактором В, є більше вираженими, ніж випадкові відмінності.

Комплект №3: гіпотези про вплив взаємодії факторів А та В

H_0 Вплив фактора А на ставлення до закону однаковий при різних градаціях фактора В, і навпаки.

H_1 Вплив фактора А на ставлення до закону різний при різних градаціях фактора В, і навпаки.

5 крок: підготовка до обчислень, обчислення окремих величин. Цей крок та сам подальший алгоритм обчислення запропоновано О. В. Сидоренко (табл. 8.7).

Таблиця 8.7.

ВЕЛИЧИНА	НАЗВА	ЗНАЧЕННЯ
T_A	суми по градаціях фактора А	$T_{A1} = 45, T_{A2} = 40$
$\sum T_A^2$	сума квадратів цих сум	$\sum T_A^2 = 45^2 + 40^2 = 3625$
T_B	суми по градаціях фактора В	$T_{B1} = 45, T_{B2} = 40$
$\sum T_B^2$	сума квадратів цих сум	$\sum T_B^2 = 45^2 + 40^2 = 3625$
T_{AB}	сума по коміркам	30, 15, 15, 25
$\sum T_{AB}^2$	сума квадратів цих сум	$\sum T_{AB}^2 = 30^2 + 15^2 + 15^2 + 25^2 = 1975$
n	кількість досліджуваних у кожній комірці	$n = 4$
a	кількість градацій фактора А	$a = 2$
b	кількість градацій фактора В	$b = 2$
N	загальна кількість індивідуальних значень	$N = 16$
$\sum X_i$	сума всіх індивідуальних значень	$\sum X_i = 85$
$\frac{1}{N}(\sum X_i)^2$	константа	$\frac{1}{N}(\sum X_i)^2 = 451,653$
$\sum X_i^2$	сума квадратів індивідуальних значень	$\sum X_i^2 = 529$

6 крок: обчислення критерію Фішера. Обчислюватися будуть такі величини (табл. 8.8):

SS_A варіативність ознаки, обумовлена фактором А

SS_B варіативність ознаки, обумовлена фактором В

SS_{AB} варіативність ознаки, обумовлена взаємодією факторів А та В

$SS_{\text{заг}}$ загальна варіативність ознаки

$SS_{\text{вип}}$ варіативність ознаки, обумовлена випадковими факторами

df число ступенів свободи

MS середній квадрат

$F_{\text{емп}}$ емпіричне значення критерію F

$F_{\text{крит}}$ критичне значення критерію F

Таблиця 8.8

ВЕЛИЧИНА	ФОРМУЛА ДЛЯ ОБЧИСЛЕННЯ	РЕЗУЛЬТАТ
SS_A	$SS_A = \frac{1}{nb} \sum T_A^2 - \frac{1}{N} \left(\sum X_i \right)^2$	$SS_A = 1,562$
SS_B	$SS_B = \frac{1}{nb} \sum T_B^2 - \frac{1}{N} \left(\sum X_i \right)^2$	$SS_B = 1,562$
SS_{AB}	$SS_{AB} = \frac{1}{n} \sum T_{AB}^2 - \frac{1}{N} \left(\sum X_i \right)^2 - SS_A - SS_B$	$SS_{AB} = 39,063$
$SS_{\text{заг}}$	$SS_{\text{заг}} = \sum X_i^2 - \frac{1}{N} \left(\sum X_i \right)^2$	$SS_{\text{заг}} = 77,437$
$SS_{\text{вип}}$	$SS_{\text{вип}} = SS_{\text{заг}} - SS_A - SS_B - SS_{AB}$	$SS_{\text{вип}} = 35,250$
df	$df_A = a - 1; df_B = b - 1; df_{AB} = df_A df_B$ $df_{\text{заг.}} = N - 1; df_{\text{вип.}} = df_{\text{заг.}} - df_A - df_B - df_{AB}$	$df_A = df_B = 1;$ $df_{AB} = 1;$ $df_{\text{заг.}} = 15;$ $df_{\text{вип.}} = 12$
MS	$MS_A = \frac{SS_A}{df_A}; MS_B = \frac{SS_B}{df_B};$ $MS_{AB} = \frac{SS_{AB}}{df_{AB}}; MS_{\text{вип.}} = \frac{SS_{\text{вип.}}}{df_{\text{вип.}}}$	$MS_A = 1,562;$ $MS_B = 1,562;$ $MS_{AB} = 39,063;$ $MS_{\text{вип.}} = 2,938$
$F_{\text{емп}}$	$F_{\text{емп}(A)} = \frac{MS_A}{MS_{\text{вип.}}}; F_{\text{емп}(B)} = \frac{MS_B}{MS_{\text{вип.}}};$ $F_{\text{емп}(AB)} = \frac{MS_{AB}}{MS_{\text{вип.}}}$	$F_{\text{емп}(A)} = 0,532;$ $F_{\text{емп}(B)} = 0,532;$ $F_{\text{емп}(AB)} = 13,296$
$df_1; df_2$	$df_1 = df_A; df_1 = df_B$ (відповідно для А і В) $df_1 = df_{AB}$ (для взаємодії А і В); $df_2 = df_{\text{вип.}}$	$df_1 = 1;$ $df_2 = 12$
$F_{\text{крит}}$	Знаходять за таблицями критичних значень	$F_{\text{крит}} = 4,75$ $(p \leq 0,05);$ $F_{\text{крит}} = 0,933$ $(p \leq 0,01)$
Перевірка гіпотез	$\begin{cases} F_{\text{емп}} < F_{\text{крит}} \Rightarrow H_0 \\ F_{\text{емп}} > F_{\text{крит}} \Rightarrow H_1 \end{cases}$	$F_{\text{емп}(A)} < F_{\text{крит}} \Rightarrow H_0$ $F_{\text{емп}(B)} < F_{\text{крит}} \Rightarrow H_0$ $F_{\text{емп}(AB)} < F_{\text{крит}} \Rightarrow H_1$

7 крок: Висновок. У комплектах гіпотез 1 та 2 приймається гіпотеза H_0 , тобто відмінності у зміні ставлення до закону, обумовлені фактором А (В), є не більше вираженими, ніж випадкові відмінності. У комплекті гіпотез 3 приймається гіпотеза H_1 про значимість впливу взаємодії двох факторів (вибору та винагороди) на зміну ставлення до закону про обмеження самоврядування студентів.

СПИСОК РЕКОМЕНДОВАНОЇ ЛІТЕРАТУРИ

1. Климчук В. О. Математичні методи у психології : навч. посіб. для студ. вищ. навч. закл. / В. О. Климчук. – Київ : Освіта України, 2009. – 288 с.
2. Москальов І. О. Застосування методів математичної статистики у психолого-педагогічних дослідженнях : навч. посіб. / І. О. Москальов, Д. П. Лисенко. – Київ : НУОУ, 2023. – 187 с.
3. Вдовенко В. В. Математичні методи в психології : навч. посіб. / В. В. Вдовенко. – Кіровоград: ПП «Центр оперативної поліграфії» Авангард», 2017. – 112 с.
4. Руська Р. В. Теорія імовірності та математична статистика в психології : навч. посіб. / Р. В. Руська. – Тернопіль, 2020. – 112 с.
5. Телейко А. Б. Математико-статистичні методи в соціології та психології : навч. посіб. / А. Б. Телейко, Р. К. Чорней. – Київ : МАУП, 2007. – 424 с.

Електронне навчальне видання

ВОРОНОВСЬКА Лариса Петрівна

**МАТЕМАТИЧНІ МЕТОДИ В ПСИХОЛОГІЧНИХ
ДОСЛІДЖЕННЯХ**

Модуль 1

КОНСПЕКТ ЛЕКЦІЙ

*(для здобувачів першого (бакалаврського) рівня вищої освіти
денної та заочної форм навчання зі спеціальності
053 – Психологія)*

Відповідальний за випуск *Л. Б. Коваленко*

За авторською редакцією

Комп'ютерне верстання *Л. П. Вороновська*

План 2024, поз. 49Л

Підп. до друку 27.02.2025. Формат 60 × 84/16.
Ум. друк арк. 7,3.

Видавець і виготовлювач:
Харківський національний університет
міського господарства імені О. М. Бекетова,
вул. Чорноглазівська (Маршала Бажанова), 17, Харків, 61002.
Електронна адреса: office@kname.edu.ua
Свідоцтво суб'єкта видавничої справи:
ДК № 5328 від 11.04.2017.