

Міністерство освіти і науки України  
Національний університет «Одеська юридична академія»

Кафедра психології

**А.О. Татьянчиков**

**КОНСПЕКТ ЛЕКЦІЙ**  
**з дисципліни**  
**МАТЕМАТИЧНІ МЕТОДИ ПСИХОЛОГІЇ**

**для здобувачів вищої освіти 4 курсу за першим (бакалаврським) рівнем**  
**денної/заочної форми навчання**  
**спеціальності «053 Психологія»**

Одеса – 2020

## Зміст

|                                                                 |    |
|-----------------------------------------------------------------|----|
| Лекція № 1. Вимірювання в психології та види шкал .....         | 3  |
| Лекція № 2. Математичні показники вибірки .....                 | 10 |
| Лекція № 3. Методи автоматизованого збору емпіричних даних..... | 23 |
| Лекція № 4. Статистичні гіпотези. Кореляційний аналіз.....      | 29 |
| Лекція № 5. Параметричні критерії розбіжностей.....             | 41 |
| Лекція № 6. Непараметричні критерії розбіжностей.....           | 53 |
| Лекція № 7. Регресійний аналіз .....                            | 64 |
| Лекція № 8. Факторний аналіз .....                              | 66 |
| Лекція № 9. Кластерний аналіз .....                             | 72 |
| Список рекомендованої літератури.....                           | 78 |

## Лекція №1

### Тема. Вимірювання в психології та види шкал

#### План

1. Математика та психологія.
2. Поняття вимірювання.
3. Види вимірювальних шкал
  - 3.1. Номінативна шкала
  - 3.2. Порядкова шкала
  - 3.3. Шкала інтервалів
  - 3.4. Шкала відношень

#### **1. Математика та психологія.**

Більш ніж 200 років тому видатний німецький філософ Іммануїл Кант із притаманною йому переконливістю обґрунтовував неспроможність психології як науки, виходячи з того, що психологічні явища не підлягають вимірюванню, а отже до них не застосовуються математичні методи. Його співвітчизник Іоганн Гербарт у своїй книзі «Психологія як наука, заново обґрунтована на досліді, метафізиці та математиці» висловлює свою думку про зв'язок психології з математикою: «Будь-яка теорія, яка бажає бути узгоджена з дослідом, перш за все, повинна бути продовжена доти, поки не прийме кількісних означень, які з'являються в досліді та лежать в його основі. Не досягнувши цього пункту, вона висить у повітрі, піддаючись будь-якому вітру сумнівів і будучи неспроможною вступити у зв'язок з іншими, вже зміцнілими поглядами». Ідеї І. Гербарта до кінця XIX ст. втілюються у життя засновниками експериментальної психології. З тих часів можливість застосування математичних методів у психології не викликає сумнівів.

З іншого боку, багато фундаментальних психологічних теорій були створені без жодної опори на математику, наприклад: теорія діяльності О.М. Леонтьєва, теорія розвиваючого навчання В.В. Давидова, психоаналіз З. Фрейда та ін. У той же час головна відмінність галузей психологічних знань,

що використовують математичні методи, полягає в тому, що їх предмет дослідження може бути не тільки описано, але й виміряно. Можливість вимірювання того чи іншого психологічного феномену, властивості, характеристики, якості відкриває доступ для застосування методів кількісного аналізу, а отже й відповідних обчислювальних процедур.

Найбільш природнім шляхом, яким математика проникає у психологію, є *математична статистика* – розділ математики, присвячений математичним методам систематизації, обробки та використання статистичних даних для наукових та практичних висновків.

*Статистика* (від лат. status – стан справ) – галузь знань, в якій викладаються загальні питання збору, вимірювання та аналізу масових кількісних або якісних даних, вивчення кількісної сторони масових суспільних явищ у числовій формі.

*Математичні методи в психології* – навчальна дисципліна, яка вивчає особливості використання методів математичної статистики для обробки даних емпіричних (заснованих на досліді, вивченні фактів, спостереженні, експерименті) психологічних досліджень і встановлення закономірностей між явищами, що вивчаються.

Застосування математичної статистики в психології дозволяє:

- 1) доводити правильність і обґрунтованість методичних прийомів і методів, які використовуються;
- 2) строго обґрунтовувати експериментальні плани;
- 3) узагальнювати дані емпіричного дослідження;
- 4) знаходити залежності між емпіричними даними;
- 5) виявляти наявність суттєвих відмінностей між групами піддослідних (наприклад, експериментальними та контрольними);
- 6) будувати статистичні прогнози;
- 7) наочно демонструвати психологічні процеси (за допомогою графіків, схем, діаграм);
- 8) вивчати структуру та проводити класифікацію процесів або явищ;

9) уникати змістовних і логічних помилок та ін.

Однак, необхідно чітко розуміти, що сама по собі математична статистика – це тільки інструментарій, який допомагає психологу ефективно розбиратися та орієнтуватися в складному емпіричному матеріалі. Найбільш важливим у будь-якому дослідженні є чітка постановка задачі, ретельне планування дослідження, побудова несуперечливих гіпотез. Крім того, найважливішим етапом дослідження з використанням методів математичної статистики є правильна інтерпретація (тлумачення, пояснення) отриманих результатів, формулювання правильних висновків.

## **2. Поняття вимірювання.**

У своїй роботі психолог достатньо часто стикається з вимірюванням індивідуально-особистісних особливостей, наприклад, тривожності, агресивності, самооцінки, креативності, властивостей нервової системи тощо. Для цього в психодіагностиці розроблені спеціальні процедури – психодіагностичні методики, тести. У процесі такого дослідження характеристики, які вивчаються, набувають кількісного представлення. Саме ці кількісні дані використовуються для статистичної обробки.

Таким чином, вимірювання у найбільш широкому розумінні визначається як приписування чисел об'єктам або подіям, яке здійснюється за певними правилами. Такі правила встановлюють відповідність між деякими властивостями об'єктів, що вивчаються, з однієї сторони, та ряду чисел – з іншої. Отже, *вимірювання* – це процедура, за допомогою якої об'єкт, що вимірюється, порівнюється з певним еталоном і отримує чисельне вираження у певному масштабі або шкалі. Саме закодована в числовій формі інформація дозволяє використовувати математичні методи та виявляти те, що без звернення до числової інтерпретації могло б залишитися прихованим. Числове представлення об'єктів або явищ дозволяє оперувати складними поняттями у більш скороченій формі. У цілому науково-дослідну роботу психолога можна проілюструвати схемою:

*Дослідник (психолог) → предмет дослідження (психічні властивості, процеси, функції) → піддослідні → дані дослідження (числові коди) → статистична обробка даних → результати статистичної обробки даних (числові коди) → інтерпретація отриманих результатів → висновки (звіт, стаття, диплом) → отримувач наукової інформації (науковий керівник, замовник, читач).*

Вивчаючи психіку людини вимірюють не її саму, а її певні психологічні особливості: риси особистості, інтелект, характеристики пізнавальної сфери тощо. Всі особливості, які можна виміряти називають змінними.

*Змінна* – властивість, яка може змінювати своє значення. Отже всі показники тестів у психології є змінними.

### **3. Види вимірювальних шкал.**

Будь-який вид вимірювання передбачає наявність одиниць вимірювання. Саме одиниця вимірювання є умовним еталоном для здійснення тих чи інших вимірювальних процедур, про який йшлося вище. Психологічні змінні здебільшого не мають власних одиниць вимірювання, тому в більшості випадків значення психологічної ознаки визначається за допомогою спеціальних вимірювальних шкал.

Згідно класифікації С. Стівенса існує чотири типи вимірювальних шкал:

- |                                                 |   |                  |
|-------------------------------------------------|---|------------------|
| 1) номінативна (номінальна, шкала найменувань); | } | неметричні шкали |
| 2) порядкова (ординарна, рангова шкала);        |   |                  |
| 3) інтервальна (шкала рівних інтервалів);       | } | метричні шкали   |
| 4) шкала відношень.                             |   |                  |

Неметричні шкали – шкали, значення яких виражені в умовних одиницях, одиниці вимірювання не можуть бути встановлені.

Метричні шкали – шкали, значення яких виражені в абсолютних числах, такі шкали мають нульовий елемент і одиниці вимірювання.

З кожною шкалою пов'язаний визначений діапазон допустимих математико-статистичних перетворень. Вихід з межі цього діапазону призводить

до того, що отримані результати позбавляються сенсу. Наприклад, до даних виражених у шкалі найменувань неможна застосовувати методи, допустимі для шкал відношень та інтервалів.

### **3.1. Номінативна шкала.**

Вимірювання в номінативній шкалі полягає в присвоєнні якій-небудь властивості або ознаці певного позначення або символу. Вимірювання зводиться до класифікації властивостей, групуванню об'єктів, об'єднанню їх у класи. Номінальна шкала визначає, що різні властивості або ознаки відрізняються одне від одного, але не передбачає кількісних операцій над ними. Наприклад: групування за такими ознаками як стать, громадянство, місце проживання, екстраверт / інтроверт, сімейний стан, відповідь «Так» / «Ні», рівень самооцінки (занижена, адекватна, завищена), рівень прояву певної ознаки (низький, середній, високий) і т.ін.

Найпростіша номінативна шкала називається дихотомічною, в якій всі об'єкти розподіляються на два класи, що не перетинаються. Дані в такій шкалі можна кодувати двома символами, наприклад: 0 та 1, А та Б, Ч та Ж і т.д.

До результатів вимірювання, представлених у номінативній шкалі, можна застосувати лише невелику кількість статистичних методів, що значно обмежує її використання. Єдиною математичною операцією є підрахунок і порівняння кількості спостережень у кожній категорії змінної.

### **3.2. Порядкова шкала.**

Порядкова шкала дає змогу не тільки класифікувати об'єкти за категоріями, а й ранжувати їх залежно від величини вимірюваної властивості, тобто упорядковувати за значенням. Для змінних, що вимірюються у цій шкалі, можливо фіксувати той факт, що одне спостереження вище або нижче, більше або менше, ніж інше. Наприклад: соціальна відповідальність особи (6 – дуже відповідальна, 5 – висока відповідальність і т.ін.); рівень успішності навчання (4 – високий, 3 – достатній, 2 – середній, 1 – низький).

Кожне окреме значення порядкової шкали характеризується лише відносно деякого іншого значення. Можна відрізнити високі показники від низьких, проте неможливо описати відмінності між ними в точних одиницях.

Для отримання значень порядкової шкали (рангів) зазвичай виконується процедура ранжування – приписування отриманим результатам порядкових чисел – рангів, в залежності від важливості або степеня вираження того чи іншого показника.

Наприклад: впорядкувати якості особистості в залежності від їх вираженості.

| <b>Я<br/>реальне</b> | <b>Якість<br/>особистості</b> | <b>Я -<br/>ідеальне</b> |
|----------------------|-------------------------------|-------------------------|
| 7                    | Відповідальність              | 1                       |
| 1                    | Комунікабельність             | 5                       |
| 3                    | Наполегливість                | 7                       |
| 2                    | Енергійність                  | 6                       |
| 5                    | Життєрадісність               | 4                       |
| 4                    | Терпимість                    | 3                       |
| 6                    | Рішучість                     | 2                       |
| Σ 28                 |                               | Σ 28                    |

Перевірка правильності ранжування здійснюється за формулою:

$$\sum R = 1 + 2 + 3 + \dots + N = \frac{N(N+1)}{2}, \text{ де } N - \text{кількість ознак, що ранжуються.}$$

Приклад 2: За деякою методикою отримано бали: 24, 25, 27, 13, 12. Необхідно проранжувати отримані результати.

### **3.3. Інтервальна шкала.**

У шкалі інтервалів кожне із можливих значень вимірних величин розташовується від найближчого на рівній відстані. Головне поняття цієї шкали – інтервал, який можна визначити як долю або частину вимірної властивості між двома сусідніми позиціями на шкалі. Дана шкала дозволяє точно визначити відмінності між категоріями або значеннями. Для інтервального вимірювання



встановлюють одиницю вимірювання (градус за Цельсієм, бал досягнень, грам і т.д.). Важливою особистістю інтервальних вимірювань є те, що властивість об'єкта не зникає, якщо в результаті вимірювання змінна набуває нульового значення. Для інтервальної шкали розташування нульової точки довільне. Інтервальні змінні завжди вимірюють в одиницях, які мають рівні інтервали. Наприклад: температура, календар, шкали з від'ємними значеннями.

### **3.4. Шкала відношень**

Дана шкала має початок відліку – 0. Нульове значення вказує на відсутність ознаки, яку вимірюють. Для відносних змінних завжди можна визначити не тільки, на скільки змінні відрізняються між собою, а й встановити у скільки разів. Шкала відношень є найбільш інформативною шкалою, до якої можна застосувати більшість математичних і статистичних методів. У шкалі відношень можна вимірювати, наприклад, кількість балів, дохід, вагу, зріст, швидкість реакції. Саме ця шкала застосовується для вимірювань у точних науках.

## Лекція № 2. Математичні показники вибірки

### План

1. Поняття вибірки. Основні види вибірок.
2. Математичні показники вибірки.
3. Нормальний закон розподілу.

#### 1. Поняття вибірки. Основні види вибірок.

Психолог експериментатор завжди вивчає якусь певну кількість людей (*вибірку*), яка відбирається з більшої за чисельністю групи, яка називається генеральною сукупністю. Таким чином, *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. Наприклад, можна вивчати особливості адаптації всіх студентів першокурсників (генеральна сукупність) за вибіркою студентів декількох вишів в межах регіону, або декількох факультетів.

Отже, *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  $n$ . У статистиці розрізняють малу ( $n < 30$ ), середню ( $30 < n < 100$ ) та велику ( $n > 100$ ) 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* – якість, що дозволяє розповсюдити отримані на вибірці висновки на всю генеральну сукупність, тобто вибірка повинна якомога повно представляти (моделювати) генеральну сукупність. Репрезентативна вибірка – 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.

Об'єм вибірки залежить перш за все від задач дослідження. В деяких випадках доцільно та достатньо одиничних випадків, якщо ті являють інтерес для науки.

Коли метою дослідження є вивчення характеристик всієї генеральної сукупності очевидно, що більший об'єм вибірки дасть більш надійні результати. Об'єм вибірки залежить також від ступеня однорідності явища, що вивчається. Як правило, чим однорідніше явище, тим меншим може бути об'єм вибірки. Крім того, об'єм вибірки залежить від статистичних методів, які передбачається використовувати.

## 2. 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).

### 2.1. Емпіричний розподіл частот.

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

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

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

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 змінної на вибірці повинна дорівнювати об'єму вибірці (загальній кількості респондентів).

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

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

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 вище вона виглядає наступним чином:

| Оцінка | Кількість студентів<br>(абсолютна частота) | % від загальної кількості<br>(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 дві вбудовані функції СЧЁТЕСЛИ та ЧАСТОТА (знаходяться в блоці статистичних функцій). Розглянемо особливості застосування кожної з них.

Функція СЧЁТЕСЛИ підраховує кількість однакових значень в заданому діапазоні клітинок. Наприклад, СЧЁТЕСЛИ(A1:A10; 5) підрахує кількість клітинок в діапазоні A1:A10, в яких записано число 5. СЧЁТЕСЛИ (A1:A10; «низький») підрахує кількість клітинок в діапазоні A1:A10, 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).

Розглянутий приклад з оцінками студентів в Excel буде мати вигляд в таблиці частот відображено формули:

|    | A          | B      | C | D              | E                   | F           |
|----|------------|--------|---|----------------|---------------------|-------------|
| 1  | № студента | Оцінка |   | Таблиця частот |                     |             |
| 2  | 1          | 5      |   | Оцінка         | Кількість студентів | %           |
| 3  | 2          | 5      |   | 5              | =СЧЁТЕСЛИ(B2:B11;5) | =E3/10*100% |
| 4  | 3          | 4      |   | 4              | =СЧЁТЕСЛИ(B2:B11;4) | =E4/10*100% |
| 5  | 4          | 5      |   | 3              | =СЧЁТЕСЛИ(B2:B11;3) | =E5/10*100% |
| 6  | 5          | 3      |   | 2              | =СЧЁТЕСЛИ(B2:B11;2) | =E6/10*100% |
| 7  | 6          | 4      |   | 1              | =СЧЁТЕСЛИ(B2:B11;1) | =E7/10*100% |
| 8  | 7          | 3      |   |                |                     |             |
| 9  | 8          | 5      |   |                |                     |             |
| 10 | 9          | 2      |   |                |                     |             |
| 11 | 10         | 4      |   |                |                     |             |
| 12 |            |        |   |                |                     |             |

Після виконання формул результати мають вигляд:

|    | A          | B      | C | D                                                                                                    | E            | F     | G |
|----|------------|--------|---|------------------------------------------------------------------------------------------------------|--------------|-------|---|
| 1  | № студента | Оцінка |   |  <b>Види частот</b> |              |       |   |
| 2  | 1          | 5      |   | Оцінка                                                                                               | Кількість ст | %     |   |
| 3  | 2          | 5      |   | 5                                                                                                    | 4            | 40,0% |   |
| 4  | 3          | 4      |   | 4                                                                                                    | 3            | 30,0% |   |
| 5  | 4          | 5      |   | 3                                                                                                    | 2            | 20,0% |   |
| 6  | 5          | 3      |   | 2                                                                                                    | 1            | 10,0% |   |
| 7  | 6          | 4      |   | 1                                                                                                    | 0            | 0,0%  |   |
| 8  | 7          | 3      |   |                                                                                                      |              |       |   |
| 9  | 8          | 5      |   |                                                                                                      |              |       |   |
| 10 | 9          | 2      |   |                                                                                                      |              |       |   |
| 11 | 10         | 4      |   |                                                                                                      |              |       |   |

В якості критерію функції СЧЁТЕСЛИ 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» (<), «більше або дорівнює» (>=), «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.

Наприклад, формула

=СЧЁТЕСЛИ(A1:A30; "=5") підрахунок комірок заданого діапазону, які містять цифру 5;

=СЧЁТЕСЛИ(A1:A30; "<>5") підрахунок комірок заданого діапазону, які НЕ містять цифру 5;

=СЧЁТЕСЛИ(A1:A30; ">5") підрахунок комірок заданого діапазону, які містять число більше за 5;

=СЧЁТЕСЛИ(A1:A30; "так") підрахунок комірок заданого діапазону, які містять слово «так»;

=СЧЁТЕСЛИ(A1:A30; "<>так") підрахунок комірок заданого діапазону, значення яких відрізняється від слова «так».

Якщо критеріїв підрахунку декілька, або ж в якості критеріїв використовується певний діапазон числових значень, то використовується функція СЧЁТЕСЛИМН(діапазон1; критерій1; діапазон2; критерій2; ...).

Наприклад:

=СЧЁТЕСЛИМН(A1:A30;">5";A1:A30;"<5") підрахунок кількості комірок заданого діапазону, які містять число більше за 5, але менше за 10.

Функція ЧАСТОТА розраховує кількість респондентів (спостережень), які потрапляють в заданий числовий інтервал.

Розглянемо приклад. 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:

| № студента | Бал тесту | № студента | Бал тесту |
|------------|-----------|------------|-----------|
| 1          | <b>14</b> | 6          | <b>28</b> |
| 2          | <b>26</b> | 7          | <b>27</b> |
| 3          | <b>10</b> | 8          | <b>5</b>  |
| 4          | <b>15</b> | 9          | <b>13</b> |
| 5          | <b>18</b> | 10         | <b>7</b>  |

Ключ до тесту:

0-16 б. – низький рівень;

17-25 б. – середній рівень (норма);

26-... б. – високий рівень

Функція має два аргументи – масив даних і масив інтервалів. ЧАСТОТА (масив даних; масив інтервалів). Масив даних – діапазон клітин з емпіричними тестовими балами, масив інтервалів – найбільше значення кожного виділеного ключем тесту інтервалу.

Для розглянутого прикладу масив інтервалів 16; 25, оскільки ключем виділяється три інтервали:

1. менше або дорівнює 16;



2. від 17 до 25;
3. більше або дорівнює 26.

З введеними формулами таблиця Excel має вигляд:

|    | A          | B         | C        | D | E              | F                      | G           |
|----|------------|-----------|----------|---|----------------|------------------------|-------------|
| 1  | № студента | Бал тесту | Інтервал |   | Таблиця частот |                        |             |
| 2  |            | 14        | 16       |   |                |                        |             |
| 3  |            | 26        | 25       |   | Рівень         | Кількість              | %           |
| 4  |            | 10        |          |   | низ            | =ЧАСТОТА(B2:B11;C2:C3) | =F4/10*100% |
| 5  |            | 15        |          |   | сер            | =ЧАСТОТА(B2:B11;C2:C3) | =F5/10*100% |
| 6  |            | 18        |          |   | вис            | =ЧАСТОТА(B2:B11;C2:C3) | =F6/10*100% |
| 7  |            | 28        |          |   |                |                        |             |
| 8  |            | 27        |          |   |                |                        |             |
| 9  |            | 5         |          |   |                |                        |             |
| 10 |            | 13        |          |   |                |                        |             |
| 11 |            | 7         |          |   |                |                        |             |

Слід зазначити, що функція ЧАСТОТА – це функція масиву. Тому її введення має ряд особливостей.

Так, для того щоб виконати функцію ЧАСТОТА необхідно:

1. Виділити діапазон клітинок, де будуть відображені результати розрахунку абсолютної частоти (для нашого прикладу це F4:F6).
2. Додати функцію Частота (Формули – Вставити функцію – Статистичні).
3. Виділити діапазон даних (B2:B11) та діапазон інтервалів (C2:C3).
4. Натиснути Ctrl+Shift+Enter.

Відносна частота розраховуються по аналогії з попереднім прикладом.

Після виконання формул результати мають вигляд:

|    | A          | B         | C        | D | E              | F         | G      |
|----|------------|-----------|----------|---|----------------|-----------|--------|
| 1  | № студента | Бал тесту | Інтервал |   | Таблиця частот |           |        |
| 2  |            | 14        | 16       |   |                |           |        |
| 3  |            | 26        | 25       |   | Рівень         | Кількість | %      |
| 4  |            | 10        |          |   | низ            | 6         | 60,00% |
| 5  |            | 15        |          |   | сер            | 1         | 10,00% |
| 6  |            | 18        |          |   | вис            | 3         | 30,00% |
| 7  |            | 28        |          |   |                |           |        |
| 8  |            | 27        |          |   |                |           |        |
| 9  |            | 5         |          |   |                |           |        |
| 10 |            | 13        |          |   |                |           |        |
| 11 |            | 7         |          |   |                |           |        |

## 2.2. 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 ( $M_o$ ) – значення, яке зустрічається серед емпіричних даних найчастіше.

Медіана ( $M_d$ ) – 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. Медіану можна також визначити, склавши таблицю накопичених частот.

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

Таким чином, МЦТ показують, навколо яких значень групується більшість емпіричних даних. Зазвичай в якості центру такого групування розглядають середнє арифметичне.

Для обчислення значення середнього арифметичного, моди, медіани в MS Excel використовуються відповідно вбудовані функції СРЗНАЧ, МОДА (або МОДА.ОДН), МЕДИАНА з категорії «Статистичні». При цьому в дужках необхідно зазначити діапазон комірок, в яких містяться емпіричні дані, або виділити діапазон курсором.

Наприклад: =СРЗНАЧ(A1:A30); =МОДА(A1:A30); =МЕДИАНА (A1:A30)

## 2.2. Міри мінливості.

Міри мінливості вказують на те, в якій мірі отримані результати відхиляються від «центру групування», що частіше за все призводить до визначення величини відхилення емпіричних даних від середнього.

Однією з таких мір мінливості є *дисперсія* – квадрат відхилення значення від середнього арифметичного. Дисперсія визначається за формулою:

$$\text{дисперсія} = \sum \frac{(\text{змінна} - \text{середнє})^2 \times \text{частота}}{\text{кількість значень} - 1} \quad \text{або} \quad s_x^2 = \frac{\sum (x_i - \bar{X})^2}{n - 1}$$

Розрахуємо дисперсію для нашого прикладу з оцінками студентів/

$x_i$  - оцінка певного студента;  $\bar{X} = 4$  – середнє арифметичне

| № 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                   |
|            | <b>40</b> | Сума            | 10                  |

$$\text{Отже, } s_x^2 = \frac{10}{10-1} = 1,11$$

У деяких випадках дисперсія є не дуже зручною ММ. Саме тому замість дисперсії досить часто використовують квадратний корінь з неї. Така величина

називається *стандартним відхиленням*. Стандартне відхилення позначається літерою «сігма» і знаходиться за формулою  $\sigma = \sqrt{s_x^2}$ .

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

Так, для розглянутого прикладу обчислення мір центральної тенденції та мір мінливості буде мати вигляд:

|    | А                     | В                      |
|----|-----------------------|------------------------|
| 1  | № студента            | Оцінка                 |
| 2  | 1                     | 5                      |
| 3  | 2                     | 5                      |
| 4  | 3                     | 4                      |
| 5  | 4                     | 5                      |
| 6  | 5                     | 3                      |
| 7  | 6                     | 4                      |
| 8  | 7                     | 3                      |
| 9  | 8                     | 5                      |
| 10 | 9                     | 2                      |
| 11 | 10                    | 4                      |
| 12 |                       |                        |
| 13 |                       |                        |
| 14 |                       |                        |
| 15 | Мода                  | =МОДА.ОДН(В2:В11)      |
| 16 | Медіана               | =МЕДИАНА(В2:В11)       |
| 17 | Середнє арифметичне   | =СРЗНАЧ(В2:В11)        |
| 18 | Дисперсія             | =ДИСП.В(В2:В11)        |
| 19 | Стандартне відхилення | =СТАНДОТКЛОН.В(В2:В11) |
| 20 |                       |                        |

Після виконання формул результати мають вигляд:

|    | A                     | B      | C    |
|----|-----------------------|--------|------|
| 1  | № студента            | Оцінка |      |
| 2  |                       | 1      | 5    |
| 3  |                       | 2      | 5    |
| 4  |                       | 3      | 4    |
| 5  |                       | 4      | 5    |
| 6  |                       | 5      | 3    |
| 7  |                       | 6      | 4    |
| 8  |                       | 7      | 3    |
| 9  |                       | 8      | 5    |
| 10 |                       | 9      | 2    |
| 11 |                       | 10     | 4    |
| 12 |                       |        |      |
| 13 |                       |        |      |
| 14 |                       |        |      |
| 15 | Мода                  |        | 5    |
| 16 | Медіана               |        | 4    |
| 17 | Середнє арифметичне   |        | 4    |
| 18 | Дисперсія             |        | 1,11 |
| 19 | Стандартне відхилення |        | 1,05 |
| 20 |                       |        |      |

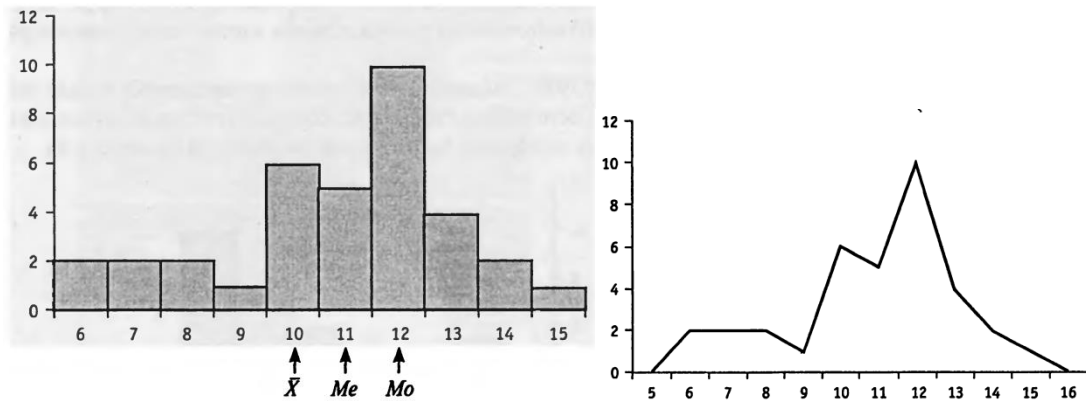
## 2.4. Степінь свободи

Число степенів свободи – це число одиниць, які вільно варіюються у складі вибірки. Так, якщо вся вибірка складається з  $n$  елементів і характеризується середнім  $\bar{X}$ , то будь-який елемент цієї сукупності може бути отриманий як різниця між величиною  $n\bar{X}$  та сумою всіх інших елементів, крім шуканого. Звідси випливає, що один елемент вибірки не має свободи варіації і завжди може бути виражений через інші. Отже, число свободи у вибіркового ряду буде дорівнювати:  $k = n - 1$ . У загальному випадку для таблиці емпіричних даних число степенів свободи буде визначатись за формулою:  $\nu = (c - 1)(n - 1)$ , де  $n$  – кількість рядків таблиці, а  $c$  – кількість її стовбців.

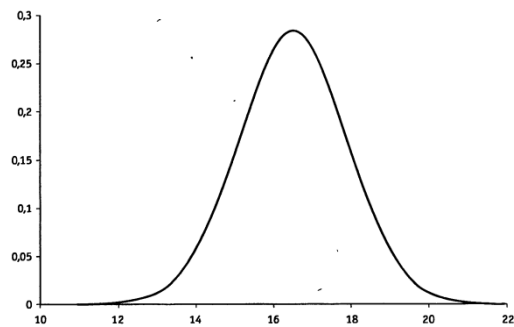
## 3. Нормальний закон розподілу.

У багатьох випадках для аналізу наявних результатів використовується їх графічне представлення.

На основі абсолютних частот емпіричних даних на вибірці можна побудувати діаграму (гістограму) та графік розподілу частот.



Якщо експериментальних даних (спостережень) достатньо багато, то в більшості випадків розподіл частот починає наближатися до вигляду симетричної кривої, схожої на дзвін.



У математичній статистиці доводиться, що по мірі збільшення числа даних такого роду розподіл наближується до вигляду, який називається кривою нормального розподілу.

Властивості нормального розподілу:

1. МЦТ мода, медіана та середнє співпадають;
2. На практиці розподіл обмежується величинами  $\bar{X} \pm 3\sigma$
3. Долю площі під кривою можна розуміти як ймовірність того, що результати будуть розташовуватися в певному діапазоні своїх значень.

## Лекція № 3

### Тема. Методи автоматизованого збору емпіричних даних

#### План

1. Коротка характеристика сервісу Google Forms.
2. Етапи створення форми з анкетною або тестом.
3. Обробка результатів.

#### **1. Коротка характеристика сервісу Google Forms.**

*Google Forms (Гугл форми)* – сервіс для створення анкет, опитувань та тестів, результати яких можна збирати через мережу Інтернет.

Переваги опитувань, створених через сервіс Google Forms:

- Легкість використання;
- Інтеграція з іншими сервісами Google (Диск, Таблиці, Документи, Презентації, Ютуб);
- Збирання результатів у Google таблиці з можливістю їх подальшої автоматизованої обробки або переведення до програми MS Excel;
- Можливість охопити опитуванням велику кількість людей за невеликий проміжок часу – достатньо просто відправити посилання на електронну пошту, в месенджері або опублікувати його на інтернет-сайті чи в соціальній мережі.

До суттєвих недоліків застосування Google Forms для психологічного тестування можна віднести відсутність вбудованих функцій та інструментів для обчислення кількісних показників тестів у відповідності до шкал та ключів методик. У зв'язку з цим виникає необхідність застосування для цього інструментарію програми MS Excel або сервісу Google таблиці, а це вимагає від користувача відповідних знань, навичок та досвіду.

## 2. Етапи створення форми з анкетною або тестом.

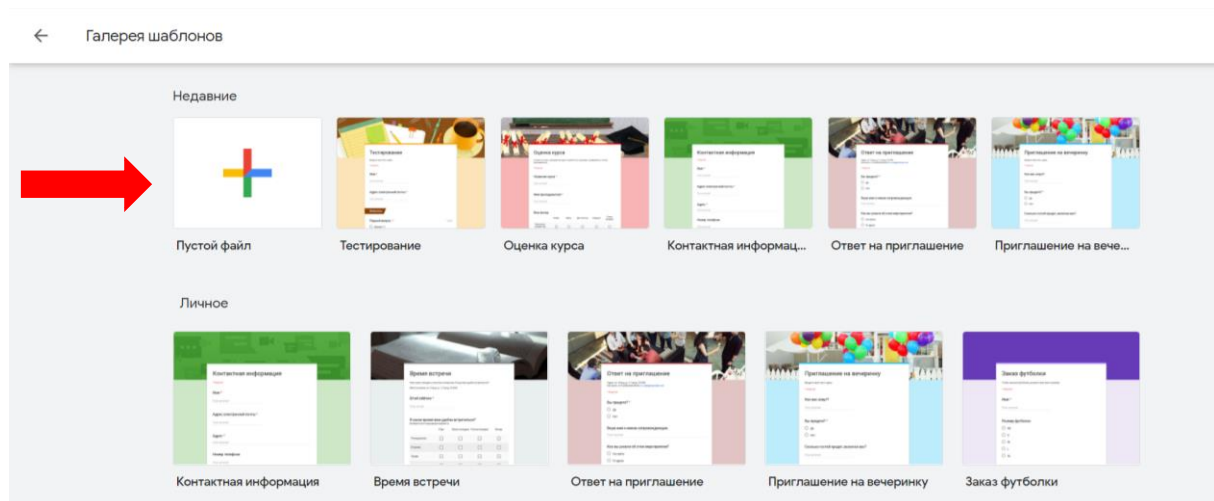
**I етап** розробки психологічного тесту в Google Forms передбачає створення форми, додавання до неї запитань тесту та варіантів відповідей.

Для реалізації цього етапу необхідно виконати наведені нижче кроки.

1. Обрати психодіагностичну методику (тест), яку необхідно перевести у Google Forms. Знайти описання цієї методики у надійному інформаційному джерелі (бажано в психодіагностичному довіднику або на перевіреному спеціалізованому інтернет-ресурсі).

2. Відкрити сервіс Google Forms за посиланням <https://docs.google.com/forms/u/0/>, або ввівши відповідний запит до пошукового Google.

3. Вибрати один з доступних шаблонів форми, або (бажано) натиснути «Пустий файл» для створення власної форми. Слід зазначити, що спеціального шаблону для психологічних тестів у сервісі не передбачено.

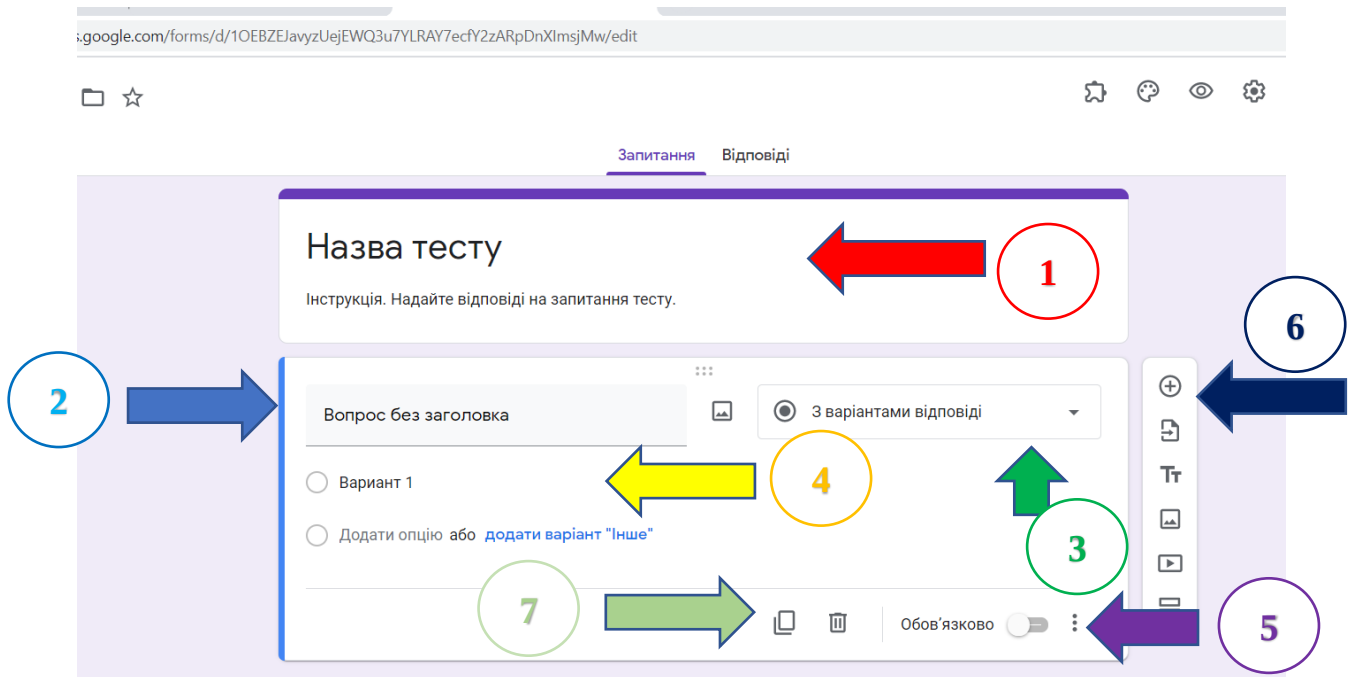


4. Введіть назву форми (наприклад, назву тесту). В поле «Опис» додайте відповідну інформацію за потреби, наприклад, інструкцію до тесту (див. малюнок, стрілка 1).

5. В поле «Питання без заголовку» введіть запитання тесту. Бажано зазначити порядковий номер запитання (2).

6. Оберіть тип варіантів відповіді, вибравши необхідний зі списку, що розкривається (3).





7. Введіть перший варіант відповіді в поле «Варіант 1» (4).
8. Щоб додати ще один варіант, натисніть «Додати опцію».
9. Якщо пропуск питання не дозволяється, ввімкніть параметр «Обов'язково» (5).
10. Для додавання наступного питання натисніть «+» (6). Якщо всі запитання мають однакові варіанти відповідей (в більшості психологічних тестів), зручно додавати наступне запитання шляхом копіювання попереднього з подальшою зміною лише формулювання запитання (7).

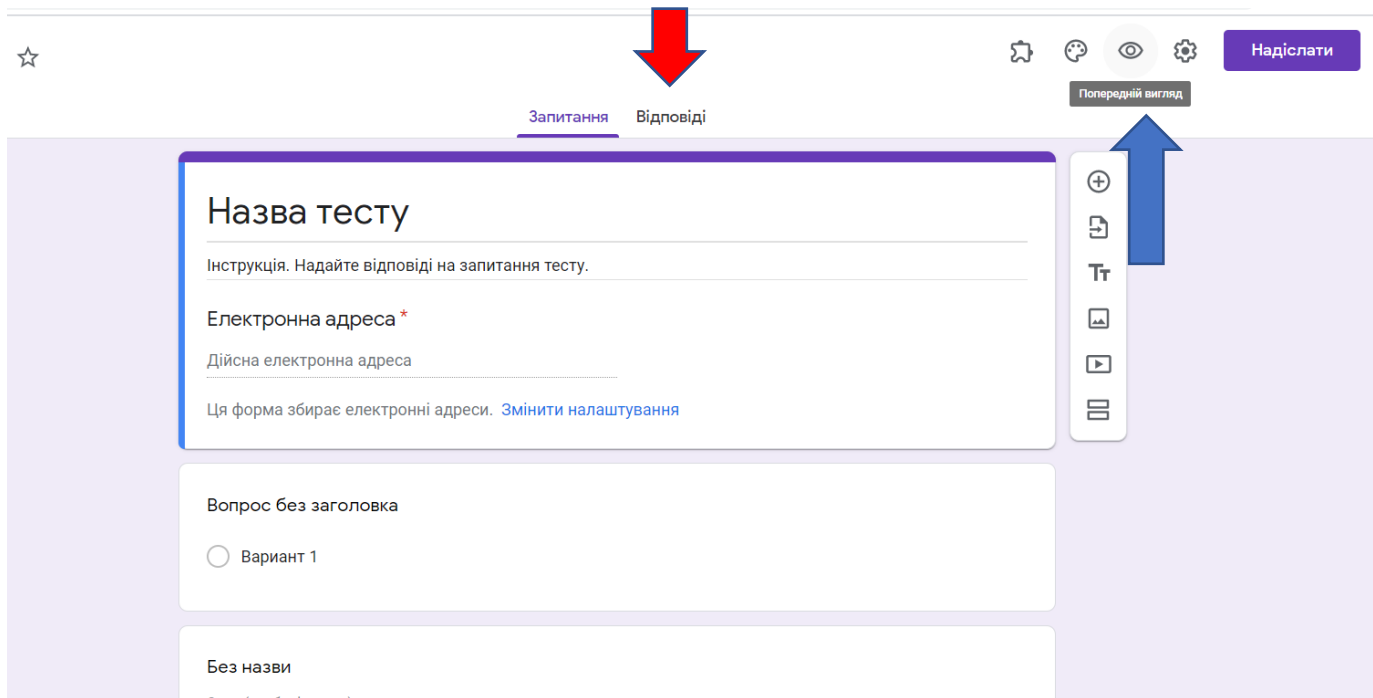
До кожного запитання можна додавати зображення та відео. Також можна імпортувати (переносити) запитання з інших форм, додавати запитання з заголовком та описом, ділити тест на окремі розділи. Управління цими параметрами відбувається через панель праворуч від форми (6).

Якщо необхідно зібрати анкетні дані респондентів (наприклад, прізвище та ім'я, дату народження, рівень освіти тощо) доцільно поділити тест на 2 розділи: перший для цих відомостей, 2 – безпосередньо для тесту.

Створений тест є анонімним, якщо адміністратор не збирає ідентифікаційні дані або не збирає адреси електронної пошти (параметр

вмикається в налаштуваннях, респонденти бачать повідомлення про це на початку тесту).

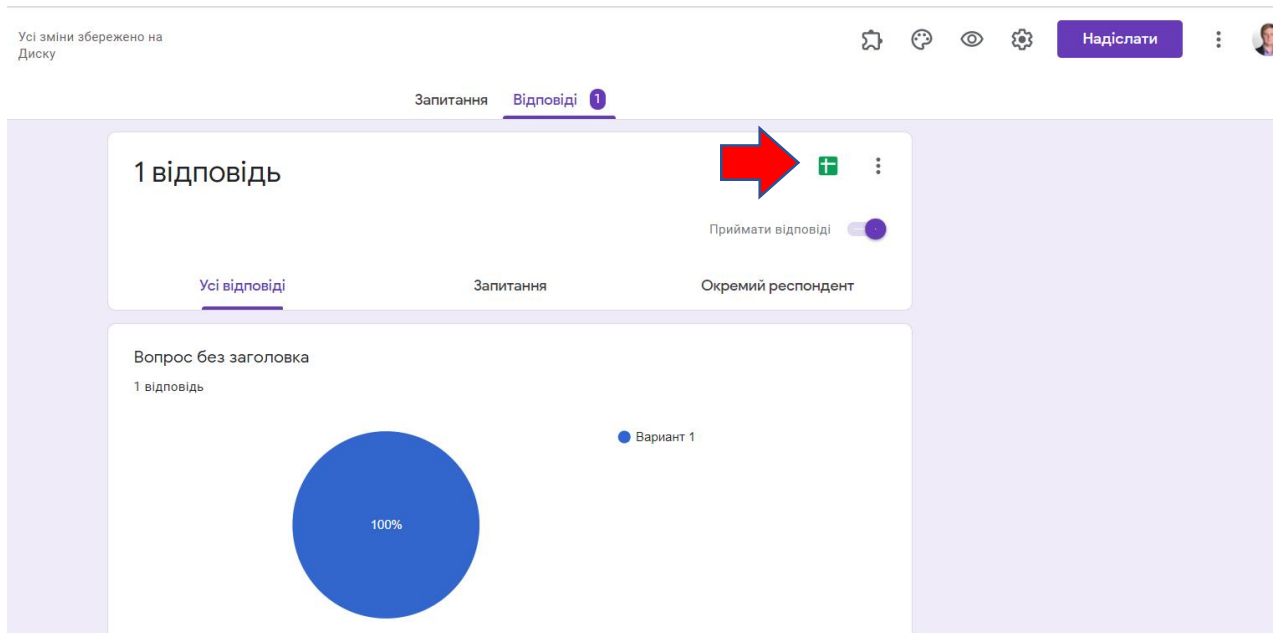
11. На будь-якому етапі створення форми її можна переглянути в тому вигляді як її побачать користувачі респонденти. Для цього необхідно натиснути на зображення ока на верхній панелі. В режимі перегляду можна також відповісти на тест, результати будуть збережені.



На цій панелі можна також відкрити налаштування дизайну форми та налаштування форми.

12. Результати опитування можна побачити, перейшовши на вкладку «Відповіді».

На даній сторінці в відображається статистика відповідей. За допомогою перемикача «Приймати відповіді» можна розпочати та зупинити збирання відповідей. Результати можна також перевести до електронної таблиці сервісу Google таблиці, натиснувши відповідну піктограму.



12. Надіслати форму респондентам можна натиснувши кнопку «Надіслати». Доступні наступні варіанти: надіслати електронною поштою, отримати посилання, вставити як фрагмент HTML-коду на власний сайт, поділитися через Facebook або Twitter.

**Зверніть увагу!** Після початку роботи з формою вона автоматично зберігається на Google Диску.

### 3. Обробка результатів.

Особливістю тесту, створеного в Google Forms є те, що до підсумкової таблиці заносяться лише варіанти відповідей у повному формулюванні, тобто дані виражені у номінативній шкалі:

[illegible]

У цьому випадку виникає проблема переведення даних відповідей у кількісні результати (бали) у відповідності до процедури, описаної в методиці.

Так, наприклад, в методиці самооцінки психічних станів Айзенка передбачається наступні варіанти відповідей на кожне запитання: «не підходить» - 0 балів, «підходить, але не дуже» - 1 бал, «підходить» - 2 бали.

Для позбавлення від монотонної роботи з ручного переведення відповідей у бали можна виділити три можливі методи автоматизованої обробки.

### **Метод 1.**

Використовування вбудованої функції в Excel «ЕСЛИМН». Її дія полягає в наданні клітинками, які містять певний текст відповідного числового значення. Для цього необхідно на аркуші Excel скласти ще одну таблицю, аналогічну таблиці з варіантами відповідей, отриманої в результаті збору даних форми.

В першу клітинку необхідно занести формулу:

=ЕСЛИМН(B2;"Подходит";2; "Подходит, но не очень";1; "Не подходит";0)

Якщо в клітинці текст, то; якщо, то; якщо, то

## Лекція № 4

### Тема. Статистичні гіпотези. Кореляційний аналіз

#### План

1. Поняття статистичної гіпотези.
2. Поняття кореляційного аналізу. Порядок здійснення кореляційного аналізу.
3. Лінійна кореляція Пірсона.
4. Рангова кореляція Спірмена.
5. Рангова кореляція Кендалла.

#### 1. Поняття статистичної гіпотези.

Отримані в результаті емпіричного дослідження на певній вибірці дані є підставою судження про генеральну сукупність. Однак, в силу дії випадкових ймовірнісних причин, оцінка параметрів генеральної сукупності завжди буде супроводжуватися похибкою. Тому такого роду оцінки повинні розглядатися як припустимі, а не остаточні твердження. Такі припущення про властивості та параметри генеральної сукупності називаються *статистичними гіпотезами*.

Статистична гіпотеза – припущення про властивості та параметри генеральної сукупності.

Сутність перевірки статистичної гіпотези полягає в тому, щоб встановити, чи узгоджується емпіричні дані та висунута гіпотеза. Таким чином, основною задачею математичної статистики є науково-обґрунтована перевірка статистичних гіпотез.

При перевірці статистичних гіпотез використовуються два поняття: так звана нульова ( $H_0$ ) та альтернативна ( $H_1$ ) гіпотеза. Прийнято вважати, що нульова гіпотеза – це гіпотеза про схожість, а альтернативна – гіпотеза про відмінність. Нульова гіпотеза виходить з того, що різниця середніх та різниця стандартних відхилень двох вибірок дорівнюють 0.

Отже,

*Нульова гіпотеза ( $H_0$ )* – це гіпотеза про схожість вибірок, відсутність значущих взаємозв'язків між змінними. Нульова гіпотеза завжди стверджує випадковий характер всіх отриманих результатів дослідження.

*Альтернативна гіпотеза ( $H_1$ )* – гіпотеза про відмінність вибірок, наявність значущих взаємозв'язків між змінними. Альтернативна гіпотеза стверджує, що отримані результати дослідження носять не випадковий характер та відображають властивості генеральної сукупності.

При обґрунтуванні статистичного висновку необхідно вирішити питання, де проходить лінія між прийняттям та відхиленням нульової гіпотези. У силу наявності у дослідженні випадкових впливів ця межа не може бути проведена абсолютно точно. Вона базується на понятті *рівня значимості*.

*Рівень значущості* – ймовірність помилкового відхилення нульової гіпотези (позначається  $P$  або  $\alpha$ ). Вважається, що нижчим рівнем статистичної значущості є рівень  $P=0,05$ , достатнім –  $P=0,01$ , вищим –  $P=0,001$ . Наприклад, рівень значущості  $P=0,05$  означає, що допускається 5 помилок на вибірку з 100 спостережень. У статистичних висновках зазвичай записують  $P \leq 0,05$  або  $P \leq 0,01$ . Рівні значущості на рівнях  $P \leq 0,05$ ,  $P \leq 0,01$ ,  $P \leq 0,001$  ще називають п'ятипроцентним, однопроцентним та 0,1-процентним рівнем значущості.

Правило прийняття статистичного висновку наступне: на основі отриманих емпіричних даних дослідник підраховує за вибраним ним статистичним методом так звану емпіричну статистику, або емпіричне значення (позначимо  $Z_{\text{емп.}}$ ).

*Емпіричне значення* статистичного критерію – числове значення, яке розраховується за спеціальною формулою на основі емпіричних даних дослідження. Емпіричне значення можна отримати за допомогою розрахунку у комп'ютерних програмах.

Крім емпіричного значення критерію існує *критичне значення* критерію – деяка постійна величина.

*Критичне значення* статистичного критерію – стале числове значення, яке не залежить від емпіричних даних дослідження та знаходиться за спеціальними довідковими таблицями критичних значень обраного статистичного критерію з урахуванням об'єму вибірки та рівня статистичної значущості.

Загальна схема прийняття або відкидання статистичної гіпотези:

- 1) розрахувати емпіричне значення статистичного критерію;
- 2) за таблицею критичних значень визначити критичні значення статистичного критерію на рівнях значущості  $P \leq 0,05$ ,  $P \leq 0,01$ ;

3) порівняти емпіричні та критичні значення статистичного критерію: якщо  $З_{емп.} < З_{кр.}$  при  $P \leq 0,05$ , то говорять про прийняття нульової гіпотези на рівні значущості  $P \leq 0,05$ . Якщо  $З_{емп.} < З_{кр.}$  при  $P \leq 0,01$ , то говорять про прийняття нульової гіпотези на рівні значущості  $P \leq 0,01$ . Якщо  $З_{емп.} > З_{кр.}$ , то приймається альтернативна гіпотеза на відповідному рівні значущості. Інтервал значущості на рівні  $P \leq 0,05$  (між двома критичними значеннями) називається також *зоною невизначеності*. Слід зазначити, що значущості критерію на даному рівні вважається достатньою для приймання альтернативної гіпотези для більшості психологічних досліджень.

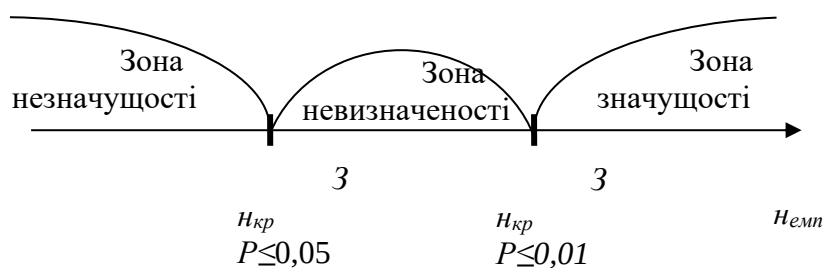


Рис. 1. Графічна схема визначення значущості статистичного критерію.

## 2. Поняття кореляційного аналізу, його призначення і види.

Досить часто психолога цікавить, як пов'язані між собою дві чи більша кількість змінних у одній або декількох групах досліджуваних. Наприклад, чи можуть учні з високим рівнем тривожності показувати стабільні навчальні

досягнення, або чи пов'язана тривалість роботи психолога у школі з його заробітною платою, або з чим більше пов'язаний рівень розумового розвитку учнів – з їх успішністю по математиці чи по літературі. Пошуком відповідей на такого роду запитання займається кореляційний аналіз, у межах якого обчислюється міра зв'язку між двома рядами значень (змінними) – коефіцієнт кореляції.

*Кореляційний аналіз* – метод математико-статистичної обробки емпіричних даних, який дозволяє встановити взаємозв'язок між змінними.

*Коефіцієнт кореляції* ( $r_{xy}$ ) – числове значення, яке обчислюється за відповідними формулами кореляційного аналізу шляхом підстановки до них емпіричних даних дослідження за змінними  $X$  та  $Y$ .

Коефіцієнт кореляції завжди лежить в межах від -1 до 1. При цьому значення 0 означає відсутність будь якого, навіть найменшого, зв'язку; додатній коефіцієнт означає прямопропорційний (прямий) зв'язок, від'ємний – зворотньопропорційний (зворотній). Чим більше значення коефіцієнту наближується до 1 або -1, тим більшим є рівень зв'язку.

*Прямопропорційний зв'язок* має місце у випадку якщо при збільшенні (зменшенні) значення змінної  $X$ , збільшується (зменшується) значення змінної  $Y$ .

Наприклад, чим більше рівень інтелекту, тим вище рівень успішності навчання та навпаки.

*Зворотньопропорційний зв'язок* має місце у випадку якщо при збільшенні (зменшенні) значення змінної  $X$ , зменшується (збільшується) значення змінної  $Y$ .

Наприклад, чим вище рівень тривожності, тим нижче рівень самооцінки та навпаки.

Перед здійсненням кореляційного аналізу формуються статистичні гіпотези. Нульова гіпотеза  $H_0$  вказує на відсутність значущого зв'язку між показниками, альтернативна  $H_1$  – про його наявність.



При цьому важливо знати, що кореляційний аналіз не встановлює причинно-наслідкового зв'язку між показниками, а лише демонструє наявність або відсутність взаємозв'язку між ними. Кореляційний аналіз не відповідає на питання який показник залежить від якого, тому кореляційний зв'язок не дорівнює кореляційній залежності.

При інтерпретації результатів кореляційного аналізу необхідно вирішити питання, де проходить лінія між прийняттям та відхиленням нульової гіпотези. У силу наявності у дослідженні випадкових впливів ця межа не може бути проведена абсолютно точно. Вона базується на понятті рівня *значимості* – ймовірності помилкового відхилення нульової гіпотези (позначається  $P$  або  $\alpha$ ). Вважається, що нижчим рівнем статистичної значимості є рівень  $P=0,05$ , достатнім –  $P=0,01$ , вищим –  $P=0,001$ . Наприклад, рівень значимості  $P=0,05$  означає, що допускається 5 помилок на вибірку з 100 елементів. У статистичних висновках зазвичай записують  $P \leq 0,05$  або  $P \leq 0,01$ .

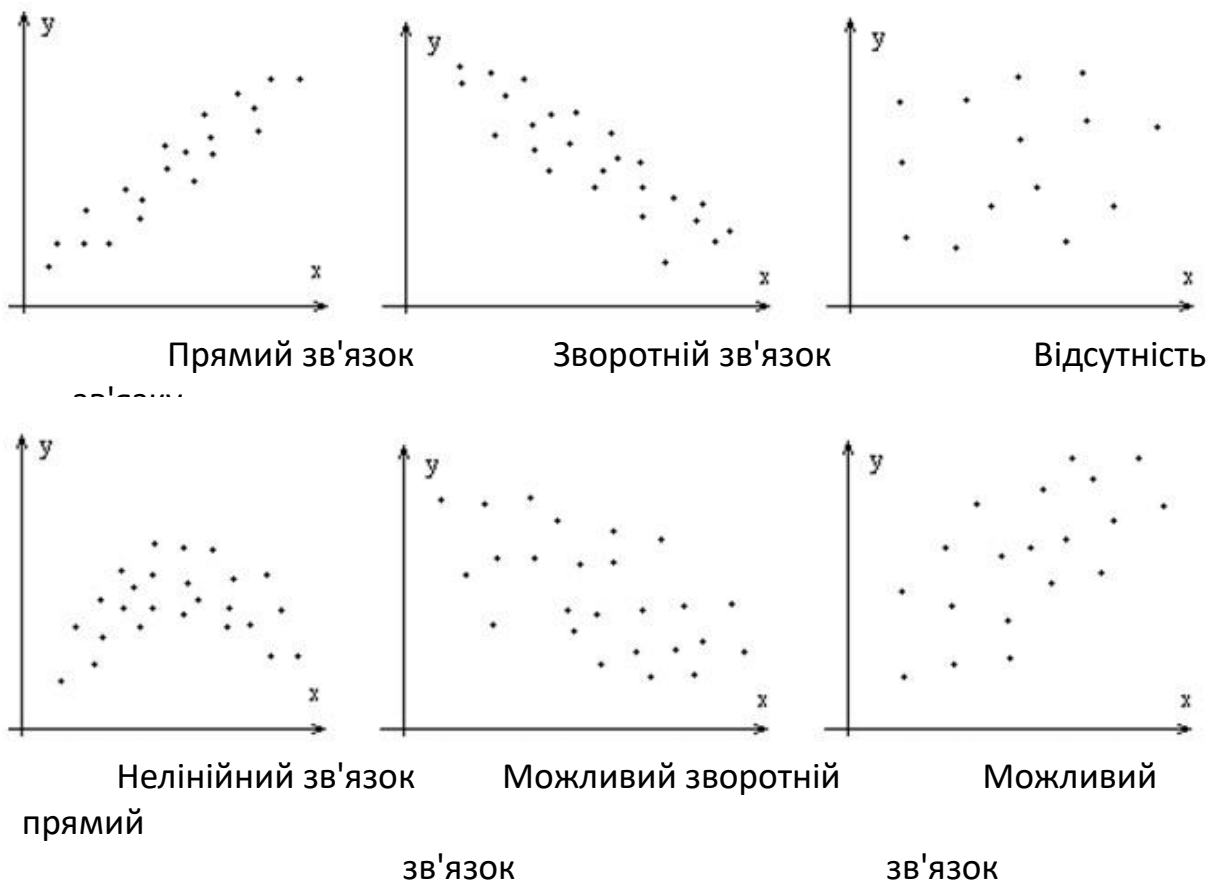
Правило прийняття статистичної гіпотези наступне: на основі отриманих емпіричних даних дослідник підраховує емпіричне значення коефіцієнту кореляції ( $r_{емп.}$ ). Надалі емпіричне значення порівнюється з двома критичними величинами, які відповідають рівням значимості  $P=0,05$  та  $P=0,01$  (позначимо  $r_{кр.}$ ), які знаходять за відповідними статистичними таблицями. Емпіричне значення порівнюється з критичними. Якщо  $r_{емп.} < r_{кр.}$  при  $P \leq 0,05$ , то говорять про прийняття нульової гіпотези на рівні значущості  $P \leq 0,05$ . Якщо  $r_{емп.} < r_{кр.}$  при  $P \leq 0,01$ , то говорять про прийняття нульової гіпотези на рівні значущості  $P \leq 0,01$ . Якщо  $r_{емп.} > r_{кр.}$  то приймається альтернативна гіпотеза на відповідному рівні значущості.

Отже, при розрахунку на основі емпіричних даних отримується емпіричне значення коефіцієнту кореляції. Для визначення його значимості, його необхідно порівняти з критичним значенням – гранично допустимим, при якому зв'язок починає або перестає бути значимим. Критичне значення залежить від об'єму вибірки (чим вона більше тем меншим є значимий коефіцієнт кореляції) та знаходиться за допомогою статистичних таблиць або обчислюється за

допомогою спеціальних комп'ютерних програм на необхідному рівні значимості ( $P \leq 0,05$ ;  $P \leq 0,01$ ;  $P \leq 0,001$ ). Емпіричне значення порівнюється з критичними та оцінюється потрапляння його у зону значимості.

Для попередньої оцінки наявності або відсутності значущого кореляційного зв'язку між змінними (показниками) застосовуються так звані діаграми розсіювання – множина точок, яка відповідає емпіричним даним дослідження в демонструє взаємозв'язок між ними на координатній площині.

По осі X позначаються емпіричні дані одного показника, по осі Y – іншого. Чим більш кучна та сформована «хмара» точок, тим більш значущий зв'язок між показниками. Наприклад:



Існує три основних методи кореляційного аналізу:

- лінійна кореляція Пірсона;
- рангова кореляція Спірмена;
- рангова кореляція Кендалла.

Змінні  $X$  та  $Y$ , між якими встановлюється значущість взаємозв'язку можуть бути виміряні в різних шкалах, це визначає вибір відповідного методу кореляційного аналізу (див. табл.).

| Тип шкали                            |                                      | Мера связи                       |
|--------------------------------------|--------------------------------------|----------------------------------|
| Переменная $X$                       | Переменная $Y$                       |                                  |
| Интервальная или отношений           | Интервальная или отношений           | Коэффициент Пирсона $r_{xy}$     |
| Ранговая, интервальная или отношений | Ранговая, интервальная или отношений | Коэффициент Спирмена $\rho_{xy}$ |
| Ранговая                             | Ранговая                             | Коэффициент «Т» Кендалла         |
| Дихотомическая                       | Дихотомическая                       | Коэффициент «Ф»                  |
| Дихотомическая                       | Ранговая                             | Рангово-бисериальный $R_{rb}$    |
| Дихотомическая                       | Интервальная или отношений           | Бисериальный $R_{бис}$           |
| Интервальная                         | Ранговая                             | Не разработан                    |

Найбільш розповсюдженими є коефіцієнт лінійної кореляції Пірсона та коефіцієнт рангової кореляції Спірмена.

### 3. Лінійна кореляція Пірсона.

Коефіцієнт кореляції Пірсона розраховується за формулою:

$$r_{xy} = \frac{n \cdot \sum(x_i \cdot y_i) - (\sum x_i \cdot \sum y_i)}{\sqrt{[n \cdot \sum x_i^2 - (\sum x_i)^2] \cdot [n \cdot \sum y_i^2 - (\sum y_i)^2]}}$$

Для застосування коефіцієнту кореляції Пірсона необхідно виконувати наступні умови:

- 1) Порівнювані змінні повинні бути представлені в метричній шкалі (інтервальній або відношень);
- 2) розподіл змінних повинен бути близьким до нормального;
- 3) число ознак у порівнюваних змінних повинне бути однаковим (кількість спостережень однаковим для обох змінних).

### Порядок розрахунку лінійної кореляції в програмі Excel

- 1) Занесіть емпіричні дані для розрахунку у таблицю.
- 2) Встановіть табличний курсор у порожню комірку. Додайте до комірки вбудовану функцію обчислення коефіцієнта лінійної кореляції Пірсона (Вкладка Формули → Додати функцію → Статистичні → PEARSON).
- 3) У діалоговому вікні, що відкрилося в якості аргументів «Масив 1» та «Масив 2» додайте діапазони комірок з даними за досліджуваними змінними. Натисніть «ОК».
- 4) В активній комірці з'явиться емпіричний коефіцієнт кореляції. Округліть його до сотих (ПКМ → Формат комірок → Кількість знаків після коми).
- 5) За таблицею критичних значень коефіцієнта лінійної кореляції Пірсона (див. Додаток А) знайдіть пару критичних значень для даного об'єму вибірки (на рівнях  $P \leq 0,05$  та  $P \leq 0,01$ ). Порівняйте з ними отримане емпіричне значення коефіцієнту кореляції.

Для побудови **кореляційної матриці** необхідний встановлений пакет «Аналіз даних».

Перевірте, чи встановлений у програмі Excel на вашому комп'ютері пакет «Аналіз даних». Для цього відкрийте вкладку «Дані» та перевірте наявність блоку «Аналіз» (див. Рис. 2).

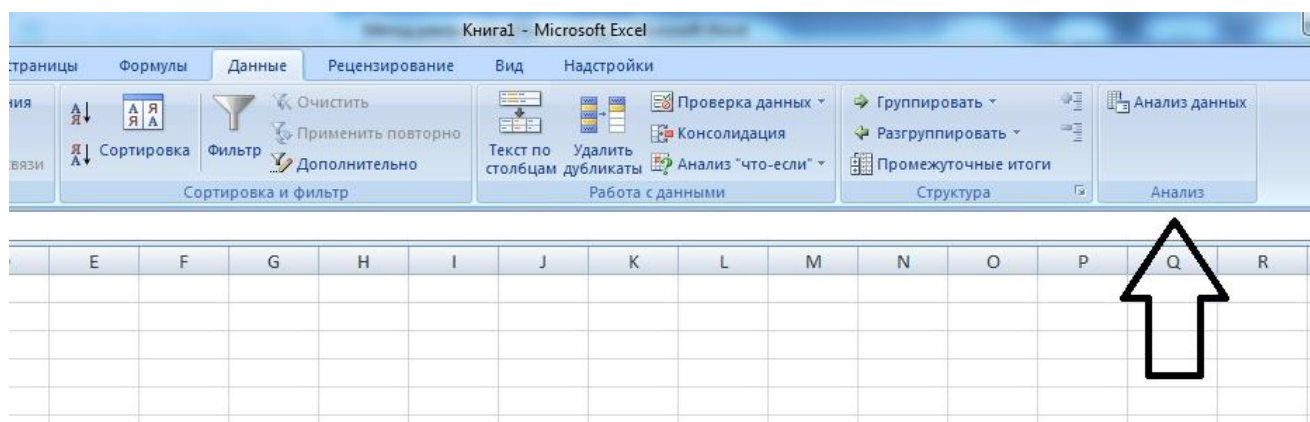


Рис. 2

Якщо пакет «Аналіз даних» не встановлений, виконайте наступні кроки:

- Відкрийте вкладку Файл, натисніть кнопку Параметри та оберіть категорію Надналаштування. Якщо ви використовуєте MS Excel 2007, натисніть кнопку Офіс , а далі – кнопку Параметри Excel.

- У списку, який розкривається Управління оберіть пункт Надлаштування для Excel.

- У діалоговому вікні, що відкрилось, встановіть параметр Пакет аналізу та натисніть ОК.

- Перейдіть на вкладку Дані та натисніть кнопку Аналіз даних блоку Аналіз (див. Рис.2).

- Оберіть інструмент аналізу Кореляція, натисніть ОК.

- У діалоговому вікні, яке відкрилось, оберіть вхідний інтервал, виділивши скопійовану вище таблицю без назв стовбців (див. Табл. 2).

- Виділіть вихідний інтервал, якому буде розташована матриця кореляції. Натисніть ОК.

- У побудованій матриці запишіть назви змінних по рядкам і стовпцям

- За таблицею критичних значень коефіцієнта лінійної кореляції Пірсона встановіть значущість отриманих у матриці коефіцієнтів і відмітьте їх наступним чином: \* - значущість на рівні  $p < 0,05$ , \*\* - значущість на рівні  $p < 0,01$ .

### **Порядок розрахунку лінійної кореляції в програмі SPSS**

- 1) Занесіть емпіричні дані для розрахунку у таблицю або скопіюйте їх з аркуша Excel.
- 2) Виберіть вкладку Аналіз даних, оберіть параметр Кореляції → Парні.
- 3) У вікні, що відкрилося, оберіть змінні, які будуть досліджуватися на взаємозв'язок.
- 4) Встановіть позначку напроти «Пірсона» та Відмічати значущі кореляції.
- 5) Натисніть ОК.

#### 4. Рангова кореляція Спірмена

Коефіцієнт рангової кореляції Спірмена належить до непараметричних показників зв'язку між змінними, які вимірюються у неметричній ранговій шкалі. При розрахунку цього коефіцієнту не потрібно жодних припущень про характер розподілу однак у генеральній сукупності. Даний коефіцієнт визначає ступінь тісноти зв'язку порядкових ознак, які в цьому випадку являють собою ранги порівняних величин.

Коефіцієнт кореляції Спірмена розраховується за формулою:

$$\rho = 1 - \frac{6 \cdot \sum (D^2)}{n \cdot (n^2 - 1)}$$

де  $n$  – кількість ознак, що ранжуються (показників, досліджуваних);  
 $D$  – різниця рангів по двом змінним для кожного досліджуваного.

За наявності однакових рангів формула розрахунку коефіцієнта кореляції Спірмена буде дещо іншою. У цьому випадку у формулу додається два нових члена, які враховують однакові ранги, які називаються поправками на однакові ранги. Вони розраховуються за формулами:

$$D_1 = \frac{n^3 - n}{12}$$

$$D_2 = \frac{k^3 - k}{12},$$

де  $n$  – число однакових рангів у першому стовбці;

$k$  – число однакових рангів у другому стовбці.

Якщо є дві групи однакових рангів у будь-якому стовбці, то формула поправки буде мати вигляд:

$$D_3 = \frac{(n^3 - n) + (k^3 - k)}{12},$$

де  $n$  – число однакових рангів у першій групі стовбця;

$k$  – число однакових рангів у другій групі стовбця.

Тоді коефіцієнт кореляції дорівнює:

$$\rho_{\text{сп}} = 1 - \frac{6 \cdot \sum d^2 + D1 + D2 + D3}{n \cdot (n^2 - 1)}.$$

Для застосування коефіцієнта кореляції Спірмена, необхідно виконувати наступні умови:

1. Порівнювані змінні повинні бути отримані у ранговій (порядковій) шкалі, але можуть бути виміряні у шкалі інтервалів чи відношень.
2. Характер розподілу величин не має значення.
3. Число ознак, які варіюються у порівнюваних змінних повинне бути однаковим.

Таблиці для визначення критичних значень коефіцієнтів Спірмена розраховані від числа ознак  $n=5$  до  $n=40$ . При більшому числі порівнюваних змінних при більшій кількості змінних слід використовувати таблицю для коефіцієнта кореляції Пірсона.

### **Порядок розрахунку рангової кореляції Спірмена в програмі Excel**

У програмі Excel відсутня вбудована функція для розрахунку коефіцієнта рангової кореляції Спірмена. У різних інформаційних джерелах зустрічаються різні, часто суперечливі способи розв'язання даної задачі. В одних зазначається, що для розрахунку коефіцієнту Спірмена використовується функція KОРРЕЛ з емпіричними даними в якості аргументів, в інших у функцію KОРРЕЛ пропонується підставляти ранги емпіричних даних. Втім, на офіційній сторінці підтримки продуктів Microsoft зазначається, що дана функція розраховує формулу лінійної кореляції Пірсона, а отже є аналогом функції PEARSON. Таким чином, найбільш правильним варіантом знаходження коефіцієнту Спірмена є розрахунок за допомогою формул, введених вручну. Приклад розрахунку коефіцієнту кореляції Спірмена таким способом для випадку різних рангів наведено у прикріпленому файлі.

### **Порядок розрахунку лінійної кореляції в програмі SPSS**

- 1) Занесіть емпіричні дані для розрахунку у таблицю або скопіюйте їх з аркуша Excel.
- 2) Виберіть вкладку Аналіз даних, оберіть параметр Кореляції → Парні.
- 3) У вікні, що відкрилося, оберіть змінні, які будуть досліджуватися на взаємозв'язок.
- 4) Встановіть позначку напроти «Спірмена» та «Відмічати значущі кореляції».
- 5) Натисніть ОК.



**Лекція № 5**  
**ТЕМА 4. Параметричні критерії розбіжностей**

План

1. Поняття статистичного критерію розбіжностей.
2. Параметричні критерії. t-критерій Стьюдента.
3. F-критерій Фішера.

**1. Поняття статистичного критерію розбіжностей.**

Однією з задач, які найчастіше зустрічаються в роботі психолога, є задача порівняння результатів обстеження якої-небудь психологічної ознаки в різних умовах вимірювання (наприклад, до та після певного впливу, обстеження в експериментальній та контрольній групах, порівняння результатів у різних групах). Крім того, часто виникає необхідність оцінити характер зміни певної психологічної ознаки в одній або декількох групах у різні періоди часу. Для розв'язання таких задач використовуються спеціальні статистичні методи – статистичні критерії розбіжностей. Існує велика кількість статистичних критеріїв розбіжностей. Вибір кожного з них залежить від типу шкали, в представлені емпіричні дані, об'єму вибірки, кількості вибірок, які необхідно порівняти, якості вибірки (зв'язані, незв'язані).

*Статистичні критерії розбіжностей* – методи математико-статистичної обробки емпіричних даних, які дозволяють встановити схожість або відмінність між значеннями змінної в двох вибірках.

Критерії розбіжностей поділяються на параметричні та непараметричні критерії.

*Параметричні критерії* ґрунтуються на нормальному розподілі генеральної сукупності або використовує параметри цієї сукупності (середнє, дисперсію та ін.).

*Непараметричні критерії* не ґрунтуються на припущенні про характер розподілу генеральної сукупності та не використовує параметри цієї сукупності.

При нормальному розподілі генеральної сукупності параметричні критерії мають більш високу потужність ніж непараметричні, тобто з більшою достовірністю спроможні відкидати нульову гіпотезу, якщо вона є помилковою.

Непараметричні критерії слід обирати при:

- 1) невідомому або відмінному від нормального характері розподілу;
- 2) малому об'єму вибірки;
- 3) нефізичній природі явищ, що вивчаються (якості, психологічні характеристики, думки і т.ін.), або при даних виміряних в неметричних шкалах.

При виборі критерію розбіжності необхідно враховувати наступні чинники:

- 1) зв'язність або незв'язність вибірки;
- 2) однорідність вибірки;
- 3) об'єм вибірки.

Для вірного вибору критерію необхідно ще раз зупинитися на понятті залежних та незалежних (зв'язних та незв'язних) вибірок.

*Незалежні* вибірки характеризуються тим, що ймовірність відбору будь-якого досліджуваного однієї вибірки не залежить від відбору будь-якого досліджуваного з іншої вибірки. Навпаки, *залежні* вибірки характеризуються тим, що кожен досліджуваний однієї вибірки поставлений у відповідність за певним критерієм досліджуваному з іншої вибірки. Отже, вибірки навиваються залежними, якщо кожному представнику з однієї вибірки є можливість поставити у відповідність єдиного представника з іншої вибірки.

Приклади залежних вибірок: одна і та сама група досліджуваних у різні періоди часу (на початку та наприкінці навчання, до та після тренінгу або якогось впливу); вибірка чоловіків та їх дружин (кожному представнику вибірки чоловіків відповідає єдиний представник з вибірки дружин); вибірки батьків та їх дітей тощо.

Приклади незалежних вибірок: дві академічні групи в університеті; учні різних класів або шкіл; чоловіки та жінки тощо.

Від правильного визначення типу вибірки залежить правильність вибору статистичного критерію.

Рекомендації щодо вибору критерію розбіжності в залежності від типу вихідних даних представлено в таблиці.

| Непараметричні критерії                                                                                                                                                                                        |                                                                                                                                                                                            | Параметричні критерії                            |                                                  |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------|--------------------------------------------------|
| Зв'язні вибірки                                                                                                                                                                                                | Незв'язні вибірки                                                                                                                                                                          | Зв'язні вибірки                                  | Незв'язні вибірки                                |
| $\chi^2$ -критерій Пірсона;<br>$\lambda$ -критерій<br>Колмогорова-Смірнова;<br>Т-критерій<br>Вілкоксона;<br>G-критерій знаків;<br>$\chi^2_r$ -критерій<br>Фрідмана;<br>L-критерій Пейджа;<br>ф-критерій Фішера | $\chi^2$ -критерій Пірсона<br>$\lambda$ -критерій<br>Колмогорова-Смірнова;<br>U-критерій Манна-Вітні;<br>Q-критерій<br>Розенбаума;<br>Н-критерій<br>Крускала-Уолліса;<br>ф-критерій Фішера | t-критерій<br>Стьюдента;<br>F-критерій<br>Фішера | t-критерій<br>Стьюдента;<br>F-критерій<br>Фішера |

## 2. Параметричні критерії. t-критерій Стьюдента.

Для порівняння вибірових середніх величин, які належать до двох сукупностей даних, і для вирішення питання про те, чи відрізняються середні значення статистично достовірно один від одного, використовують t-критерій Стьюдента або його непараметричні аналоги.

Критерій розроблений англійським математиком Вільямом С. Госсетом (1876-1937). «Стьюдент» – це псевдонім, під яким В. Госсет друкував свої роботи, працюючи в Дубліні на пивоварні Артура Гіннеса. Інший дослідник, який свого часу працював на Гіннеса, опублікував статтю, що містила конфіденційну комерційну інформацію. З того часу працівникам підприємства було заборонено публікуватися Тож В. Госсет, щоб уникнути розголосу і обійти заборону, друкував свої роботи під псевдонімом.

У вітчизняній традиції прийнято говорити про t-критерій Стьюдента. В англійській літературі і в статистичних програмах вживають термін «Т-тест».

Критерій t-Стьюдента використовується у трьох випадках:

- 1) порівняння середніх показників двох залежних вибірок (t-критерій для залежних вибірок);
- 2) порівняння середніх показників двох незалежних вибірок (t-критерій для незалежних вибірок);
- 3) порівняння середнього показника однієї вибірки із певною заданою величиною (t-критерій для однієї вибірки).

Застосувавши критерій Стюдента, ми дізнаємося, наскільки статистично значимі відмінності між двома вибірками і, відповідно, наскільки впевнено можна робити висновки про ці відмінності.

Зокрема, критерій Стюдента значно потужніший за свої непараметричні аналоги, але вимагає дотримання двох вимог: 1) відповідності числових даних параметрам нормального розподілу; 2) рівності дисперсій двох вибірок, які порівнюються між собою.

***Алгоритм застосування t-критерію Стюдента:***

1. Перевірка нормальності розподілу даних у вибірках, що порівнюються.
2. Перевірка рівності (гомогенності) дисперсій незалежних вибірок або перевірка наявності прямого кореляційного зв'язку між залежними вибірками.
3. Власне обчислення t-критерію (схема обчислення критерію різниться залежно від того, порівнюються залежні чи незалежні вибірки; однаковий чи різний обсяг вибірок).
4. Порівняння отриманого емпіричного значення t-критерію із критичним табличним значенням і висновок про підтвердження чи спростування нульової гіпотези про відсутність відмінностей.

Таким чином, *t-критерій Стюдента* спрямований на оцінку розбіжностей між величинами середніх двох вибірок, які розподілені за нормальним законом. Критерій може бути застосований для зв'язних і незв'язних вибірок, причому як однакових, так і різних за об'ємом.

T-критерій Стюдента зручно розраховувати у програмі MS Excel або її аналогах (OpenOffice Calc, Google Tables) або програмі SPSS (IBM SPSS Statistic).

У програмі Excel т-критерій Стюдента розраховується за допомогою функції ТТЕСТ (у версіях 2010 та нижче) або СТЬЮДЕНТ.ТЕСТ (у нових версіях). Функція має чотири аргументи: Масив 1, Масив 2, Хвости, Тип.

Масив 1 – емпіричні дані дослідження з першої порівнювальної вибірки.

Масив 2 – емпіричні дані дослідження з 2 порівнювальної вибірки.

Хвости – кількість хвостів розподілу. Може набувати значень: 1 – використовується односторонній розподіл, 2 – двосторонній розподіл. Односторонній розподіл обирається якщо заздалегідь відомо в яку сторону (більшу або меншу) відрізняються дані однієї з порівнюваних вибірок від іншої. Наприклад, точно відомо, що результати успішності навчання в групі 1 нижче від результатів групи 2, але не відома статистична значущість цієї відмінності. В такому т-критерій даватиме більш надійні результати у порівнянні з двостороннім критерієм. У випадку, коли напрямок відмінностей невідомий або невизначений (деякі респонденти мають вищі результати, деякі – нижчі), слід використовувати двосторонній розподіл.

Параметр «Тип» набуває наступних значень: 1 – парний тест (залежні вибірки, одна вибірка в різних умовах), 2 – двовибірковий тест з рівними дисперсіями (незалежні вибірки), 3 – двовибірковий тест з різними дисперсіями (незалежні вибірки). Отже для вибору типу 2 або 3 необхідно спочатку розрахувати та порівняти дисперсії (див. тему «Математичні показники вибірки»).

Слід звернути увагу, що функція СТЬЮДЕНТ.ТЕСТ (ТТЕСТ) розраховує не саме значення т-критерію, а ймовірність  $p$  того, що вибірки не будуть відрізнятися між собою за певною ознакою, тобто ймовірність схожості двох вибірок (ймовірність прийняття нульової гіпотези). Наприклад, для того, щоб стверджувати про наявність значущих відмінностей між двома вибірками на рівні значущості  $p < 0,05$  (альтернативні гіпотеза), отримана в результаті обчислення функції ймовірність не повинна перевищувати 0,05. Для значущості відмінностей на рівні  $p < 0,01$ , отримана ймовірність не повинна перевищувати 0,01.

### Приклад

На основі оцінок за контрольний тест з іноземної мови в двох класах з різними методами навчання, необхідно встановити чи відрізняється успішність вивчення іноземної мови в даних класах. Визначити в якому класі більш ефективна методика. Емпіричні дані наведено у таблиці.

| Клас А |        | Клас Б |        |
|--------|--------|--------|--------|
| № учня | Оцінка | № учня | Оцінка |
| 1      | 7      | 1      | 6      |
| 2      | 8      | 2      | 7      |
| 3      | 3      | 3      | 6      |
| 4      | 5      | 4      | 10     |
| 5      | 6      | 5      | 9      |
| 6      | 10     | 6      | 11     |
| 7      | 8      | 7      | 10     |
| 8      | 9      | 8      | 8      |
| 9      | 9      | 9      | 7      |
| 10     | 7      | 10     | 6      |
| 11     | 6      | 11     | 5      |
| 12     | 11     | 12     | 7      |
| 13     | 7      | 13     | 10     |
| 14     | 8      | 14     | 10     |
| 15     | 4      | 15     | 4      |
| 16     | 6      | 16     | 9      |
| 17     | 3      | 17     | 3      |
| 18     | 8      | 18     | 11     |
| 19     | 9      | 19     | 7      |
| 20     | 8      | 20     | 12     |
| 21     | 10     | 21     | 9      |
| 22     | 6      | 22     | 8      |

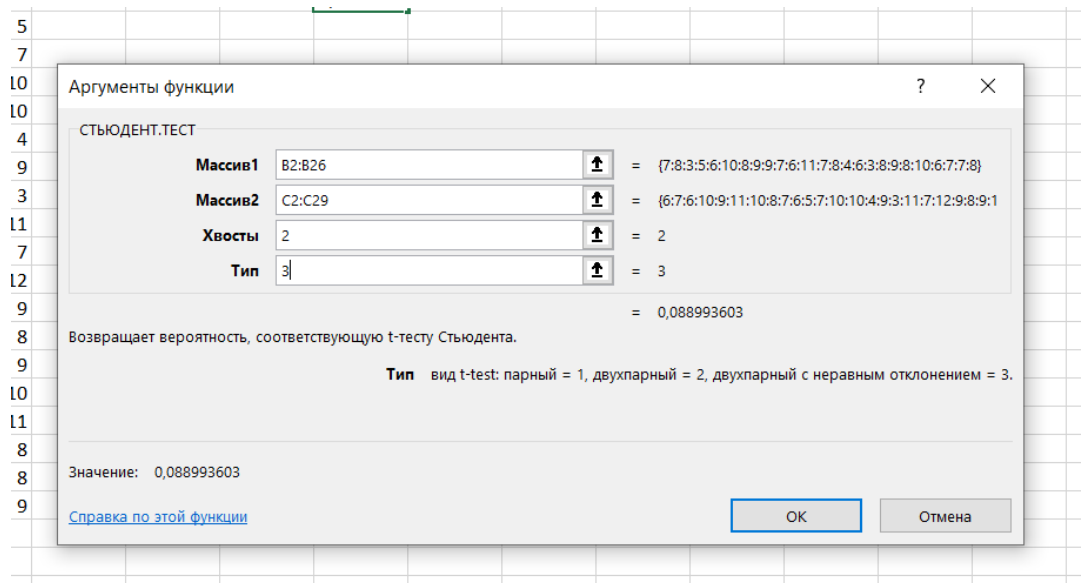
Будемо вважати, що розподіл даних в обох вибірках близький до нормального.

Альтернативна гіпотеза: «Рівень успішності вивчення іноземної мови принципово відрізняється в двох класах».

Скопіюємо дані з оцінками до таблиці Excel та знайдемо дисперсію в обох вибірках.

[illegible]

У пусту клітинку додаємо статистичну функцію СТЬЮДЕНТ.ТЕСТ. В якості аргументів Массив 1 та Массив 2 виділяємо діапазони даних двох класів. Параметр «Хвосты» ставимо «2» (двосторонній розподіл, оскільки не знаємо в якому класі результати краще). Параметр «Тип» ставимо «3» (незалежні вибірки, нерівні дисперсії).



Натискаємо ОК. Отримуємо результат 0,08899... Отже, отримана ймовірність перевищує гранично допустиму 0,05. Приймається нульова гіпотеза. Не встановлено значущих відмінностей між двома класами. Ефективність методик навчання є приблизно однаковою.

Виконаємо той самий приклад у програмі SPSS.

Створимо дві змінні «Оцінка» з числовим типом даних та «Клас» з текстовим типом даних.



|    | Имя    | Тип       | Ширина | Знаков ... | Метка | Значения | Пропущенн... | Ширина ... | Выравнивание | Мера        | Роль    |
|----|--------|-----------|--------|------------|-------|----------|--------------|------------|--------------|-------------|---------|
| 1  | Оцінка | Числовой  | 8      | 2          |       | Нет      | Нет          | 8          | По право...  | Нет данных  | Входная |
| 2  | Клас   | Текстовая | 8      | 0          |       | Нет      | Нет          | 8          | По левом...  | Номинальная | Входная |
| 3  |        |           |        |            |       |          |              |            |              |             |         |
| 4  |        |           |        |            |       |          |              |            |              |             |         |
| 5  |        |           |        |            |       |          |              |            |              |             |         |
| 6  |        |           |        |            |       |          |              |            |              |             |         |
| 7  |        |           |        |            |       |          |              |            |              |             |         |
| 8  |        |           |        |            |       |          |              |            |              |             |         |
| 9  |        |           |        |            |       |          |              |            |              |             |         |
| 10 |        |           |        |            |       |          |              |            |              |             |         |
| 11 |        |           |        |            |       |          |              |            |              |             |         |
| 12 |        |           |        |            |       |          |              |            |              |             |         |
| 13 |        |           |        |            |       |          |              |            |              |             |         |
| 14 |        |           |        |            |       |          |              |            |              |             |         |
| 15 |        |           |        |            |       |          |              |            |              |             |         |
| 16 |        |           |        |            |       |          |              |            |              |             |         |
| 17 |        |           |        |            |       |          |              |            |              |             |         |
| 18 |        |           |        |            |       |          |              |            |              |             |         |
| 19 |        |           |        |            |       |          |              |            |              |             |         |
| 20 |        |           |        |            |       |          |              |            |              |             |         |
| 21 |        |           |        |            |       |          |              |            |              |             |         |
| 22 |        |           |        |            |       |          |              |            |              |             |         |
| 23 |        |           |        |            |       |          |              |            |              |             |         |
| 24 |        |           |        |            |       |          |              |            |              |             |         |
| 25 |        |           |        |            |       |          |              |            |              |             |         |
| 26 |        |           |        |            |       |          |              |            |              |             |         |
| 27 |        |           |        |            |       |          |              |            |              |             |         |
| 28 |        |           |        |            |       |          |              |            |              |             |         |
| 29 |        |           |        |            |       |          |              |            |              |             |         |

Представление Данные    Представление Переменные

Скопіюємо данні таким чином, що оцінки одного класу опинилися під іншим в одному стовпчику. Для розрізнення двох класів між сабою візьмемо додатковий стовпчик «Клас», в якому напроти кожного респондента проставимо літеру «А» або «Б», що відповідає позначенню класу.

\*Без имени1 [Наборданных0] - Редактор данных IBM SPSS Statistics

Файл Правка Вид Данные Преобразование Анализ Прямой маркетинг Графика Сервис Окно

27 :

|    | Оцінка | Клас | пер. | пер. | пер. | пер. | пер. | пер. |
|----|--------|------|------|------|------|------|------|------|
| 10 | 7,00   | A    |      |      |      |      |      |      |
| 11 | 6,00   | A    |      |      |      |      |      |      |
| 12 | 11,00  | A    |      |      |      |      |      |      |
| 13 | 7,00   | A    |      |      |      |      |      |      |
| 14 | 8,00   | A    |      |      |      |      |      |      |
| 15 | 4,00   | A    |      |      |      |      |      |      |
| 16 | 6,00   | A    |      |      |      |      |      |      |
| 17 | 3,00   | A    |      |      |      |      |      |      |
| 18 | 8,00   | A    |      |      |      |      |      |      |
| 19 | 9,00   | A    |      |      |      |      |      |      |
| 20 | 8,00   | A    |      |      |      |      |      |      |
| 21 | 10,00  | A    |      |      |      |      |      |      |
| 22 | 6,00   | A    |      |      |      |      |      |      |
| 23 | 7,00   | A    |      |      |      |      |      |      |
| 24 | 7,00   | A    |      |      |      |      |      |      |
| 25 | 8,00   | A    |      |      |      |      |      |      |
| 26 | 6,00   | Б    |      |      |      |      |      |      |
| 27 | 7,00   | Б    |      |      |      |      |      |      |
| 28 | 6,00   | Б    |      |      |      |      |      |      |
| 29 | 10,00  | Б    |      |      |      |      |      |      |
| 30 | 9,00   | Б    |      |      |      |      |      |      |
| 31 | 11,00  | Б    |      |      |      |      |      |      |
| 32 | 10,00  | Б    |      |      |      |      |      |      |
| 33 | 8,00   | Б    |      |      |      |      |      |      |
| 34 | 7,00   | Б    |      |      |      |      |      |      |
| 35 | 6,00   | Б    |      |      |      |      |      |      |
| 36 | 5,00   | Б    |      |      |      |      |      |      |
| 37 | 7,00   | Б    |      |      |      |      |      |      |

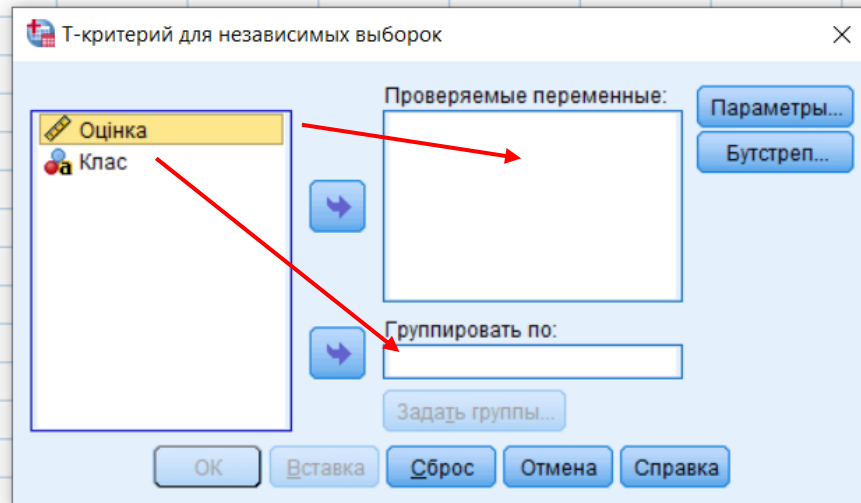
Представление Данные Представление Переменные

Для обчислення т-критерію необхідно вибрати Аналіз – Порівняння середніх та вибрати один за наступних пунктів, в залежності від умови задачі:

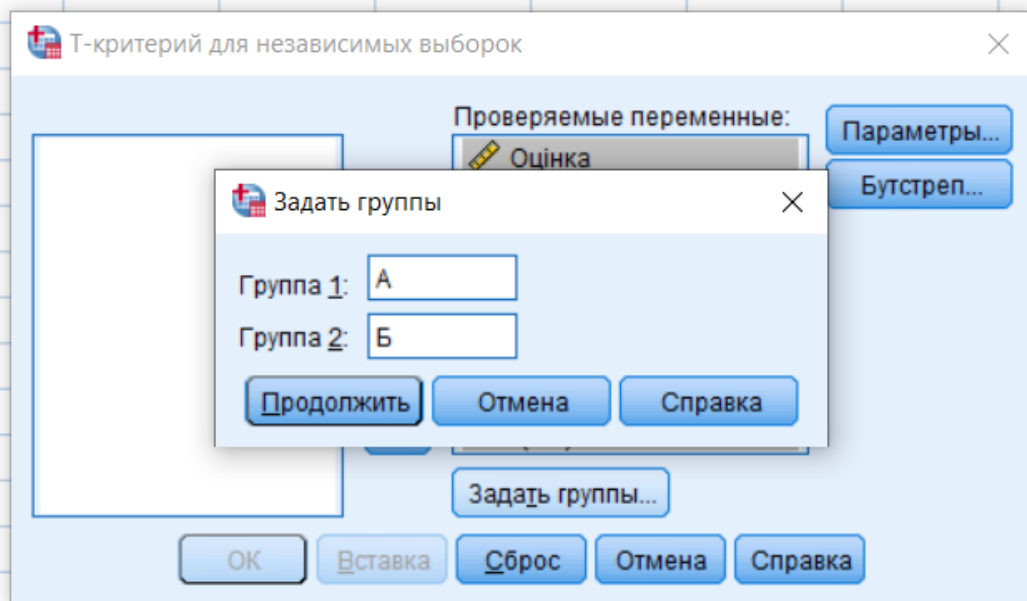
- одновибірковий т-критерій;
- т-критерій від незалежних вибірок;
- т-критерій для парних вибірок.

Для даного прикладу обрати «т-критерій від незалежних вибірок».

У вікні, що відкрилось обрати в якості змінної, яка перевіряється «Оцінка» (та змінна, за якою здійснюється аналіз), в якості змінної групування «Клас» (для розрізнення вибірок).



Натиснути кнопку «Задати групи» та ввести літери позначення груп (класів А та Б).



Натиснути «Продовжити», ОК.

Вікно виводу результату має вигляд:

[Наборданных0]

Статистика группы

|        | Клас | N  | Среднее | Среднекв. отклонение | Среднекв. ошибка среднего |
|--------|------|----|---------|----------------------|---------------------------|
| Оцінка | А    | 25 | 7,2000  | 2,04124              | ,40825                    |
|        | Б    | 28 | 8,2143  | 2,21706              | ,41898                    |

Критерий для независимых выборок

|        |                                    | Критерий равенства дисперсий Ливиня |            | t-критерий для равенства средних |        |                       |                  |                                    |                                         |         |
|--------|------------------------------------|-------------------------------------|------------|----------------------------------|--------|-----------------------|------------------|------------------------------------|-----------------------------------------|---------|
|        |                                    | F                                   | Значимость | t                                | ст.св. | Знач. (двухсторонняя) | Средняя разность | Среднеквадратичная ошибка разности | 95% доверительный интервал для разности |         |
|        |                                    |                                     |            |                                  |        |                       |                  |                                    | Нижняя                                  | Верхняя |
| Оцінка | Предполагаются равные дисперсии    | ,390                                | ,535       | -1,726                           | 51     | ,090                  | -1,01429         | ,58778                             | -2,19430                                | ,16573  |
|        | Не предполагаются равные дисперсии |                                     |            | -1,734                           | 50,945 | ,089                  | -1,01429         | ,58499                             | -2,18874                                | ,16016  |

Ймовірність нульової гіпотези (схожості вибірок) – аналог відповіді, отриманої в Excel.

Емпіричне значення критерію. Знак «-» вказує не те, що в першому вказанному класі (А) результати гірше ніж в другому (класі Б).

## Лекція № 6

### Тема. Непараметричні критерії розбіжностей

#### План

1. Поняття непараметричних критеріїв розбіжностей.
2. Хі-квадрат критерій Пірсона.
3. Кутове перетворення Фішера.

#### **1. Поняття непараметричних критеріїв розбіжностей.**

Як було зазначено в попередній лекції, статистичні критерії розбіжностей використовуються для встановлення значущості розбіжностей між двом вибірками за яким небудь показником (змінною). Наприклад, критерій допомагає відповісти на питання чи різняться між собою дві групи за рівнем інтелекту, успішністю навчання, задоволеністю професією тощо. При цьому нульова гіпотеза свідчить про схожість вибірок (відсутність значущих відмінностей), а альтернативно – про відмінність вибірок (наявність значущих відмінностей).

У попередній лекції було розглянуто параметричні критерії розбіжностей, зокрема т-критерій Стюдента.

При неможливості використання параметричних критеріїв (дані виражені в неметричних шкалах, закон розподілу ознаки відрізняється від нормального, малий об'єм вибірки) слід використовувати непараметричні критерії.

Непараметричні критерії не ґрунтуються на припущенні про характер розподілу генеральної сукупності та не використовуює параметри цієї сукупності (середнє, дисперсію, стандартне відхилення).

При нормальному розподілі генеральної сукупності параметричні критерії мають більш високу потужність ніж непараметричні, тобто з більшою достовірністю спроможні відкидати нульову гіпотезу, якщо вона є помилковою.

Отже, непараметричні критерії слід обирати при:

- 1) невідомому або відмінному від нормального характері розподілу;
- 2) малому об'єму вибірки;
- 3) нефізичній природі явищ, що вивчаються (якості, психологічні характеристики, думки і т.ін.), або при даних виміряних в неметричних шкалах.

## 2. Хі-квадрат критерій Пірсона.

Як вже було зазначено в попередній лекції, у випадку якщо емпіричні дані виражені в неметричних шкалах, якщо розподіл вибірки не є нормальним, або кількість респондентів у вибірці є недостатньою, використовуються непараметричні критерії розбіжностей (інша назва непараметричні критерії згоди).

Одним із найбільш універсальних та потужних непараметричних критеріїв розбіжностей є  $\chi^2$ -критерій Пірсона (хі-квадрат критерій).

Даний критерій дозволяє працювати з даними, отриманими в будь-якій шкалі, починаючи зі шкали найменувань та використовується у двох випадках:

- 1) розрахунок згоди емпіричного розподілу з припустимим теоретичним (гіпотеза  $H_0$  – про відсутність розбіжностей між теоретичним і емпіричним розподілами);
- 2) розрахунок однорідності двох незалежних емпіричних вибірок (гіпотеза  $H_0$  – про відсутність розбіжностей між двома емпіричними розподілами).

Для першого випадку емпіричне значення критерію розраховується за формулою:

$$\chi_e^2 = \sum_{i=1}^k \frac{(f_e - f_t)^2}{f_t}, \text{ де} \quad (1)$$

$f_e$  – емпірична частота,

$f_t$  – теоретична частота,

$k$  – кількість розрядів ознаки

Для випадку двох емпіричних розподілів додатково будується таблиця теоретичних частот ( $f_t$ ).

Для кожного значення емпіричної частоти теоретична частота знаходиться за формулою:

$$f_{tij} = \frac{\sum f_{ei} \cdot \sum f_{ej}}{\sum f_{eij}}, \text{ де} \quad (2)$$

$\sum f_{ei}$  – сума емпіричних частот за  $i$ -м рядком таблиці,

$\sum f_{ej}$  – сума емпіричних частот за  $j$ -м стовбцем таблиці,

$\sum f_{eij}$  – сума емпіричних частот за  $i$ -м рядком та  $j$ -м стовбцем таблиці.

Критичне значення критерію визначається за спеціальною таблицею з урахуванням числа ступенів свободи  $\nu = k - 1$ , де  $k$  – кількість елементів у вибірці. Якщо використовується таблиця значень, то  $\nu = (k - 1)(c - 1)$ , де  $k$  – кількість рядків,  $c$  – кількість стовбців.

Таким чином,  $\chi^2$ -критерій Пірсона – непараметричний статистичний критерій розбіжностей, який дозволяє знайти відмінності між емпіричним та теоретичним або двома емпіричними розподілами ознаки у незалежних вибірках.

Слід зазначити, що даний критерій розраховується тільки для незалежних вибірок, а теоретичні частоти не повинні бути меншими 5.

#### Приклад 1.

Психологу необхідно встановити: чи буде задоволеність працею на даному підприємстві розподілена рівномірно за наступними градаціями:

- 1 – роботою цілком задоволений;
- 2 – скоріше задоволений, чим ні;
- 3 – важко сказати, мені все одно;
- 4 – скоріше незадоволений;
- 5 – цілком незадоволений.

Для вирішення задачі була складена вибірка з 65 працівників (респондентів), в якій було проведено відповідне опитування. Отримані результати представили у вигляді таблиці частот:

| Відповідь | Кількість працівників (частота)<br>$f_{emp.}$ | Теоретична частота<br>$f_m.$ | $f_{emp.} - f_m.$ | $(f_{emp.} - f_m.)^2$ | $\frac{(f_{emp.} - f_m.)^2}{f_m.}$ |
|-----------|-----------------------------------------------|------------------------------|-------------------|-----------------------|------------------------------------|
| 1         | 8                                             | 13                           | -5                | 25                    | 1,92                               |
| 2         | 22                                            | 13                           | 9                 | 81                    | 6,23                               |

|      |    |    |    |    |                               |
|------|----|----|----|----|-------------------------------|
| 3    | 14 | 13 | 1  | 1  | 0,08                          |
| 4    | 9  | 13 | -4 | 16 | 1,23                          |
| 5    | 12 | 13 | -1 | 1  | 0,08                          |
| Сума | 65 | 65 | 0  |    | $\chi^2_{\text{емп.}} = 9,54$ |

За таблицею критичних значень знайдемо  $\chi^2_{\text{кр.}}$  для числа степенів свободи  $\nu = 4$  та двох рівнів значимості:

$$\chi^2_{\text{кр}} = \begin{cases} 9,488 & \text{для } P \leq 0,05 \\ 13,277 & \text{для } P \leq 0,01 \end{cases}$$

Таким чином  $\chi^2_{\text{емп.}} = 9,54$  потрапляє в зону невизначеності. Однак, на рівні 5% значимості можна прийняти альтернативну гіпотезу  $H_1$  про значимі відмінності теоретичного розподілу від емпіричного. Тобто задоволеність працею на даному підприємстві розподілена серед працівників нерівномірно.

На практиці набагато частіше зустрічаються задачі, в яких необхідно порівнювати два та більше емпіричних розподілів. У таких задачах за допомогою хі-квадрат критерію проводиться оцінка однорідності двох або більше незалежних вибірок і таким чином перевіряється гіпотеза про відсутність відмінностей між цими розподілами.

У найбільш простих випадках емпіричні дані двох вибірок можна представити у вигляді таблиці 2x2.

Приклад:

У двох школах психолог з'ясовував думку вчителів про організацію психологічної служби в школі. У першій школі було опитано 20 учителів, у другій – 15. Психолога цікавило питання: в якій школі психологічна служба представлена краще? Учителі надавали відповіді по номінативній шкалі: «подобається» або «не подобається». Результати опитування представлені в чотирипольній таблиці.



|              | 1 школа | 2 школа | Сума |
|--------------|---------|---------|------|
| Задоволені   | 15      | 7       | 22   |
| Незадоволені | 5       | 8       | 13   |
| Сума         | 20      | 15      | 35   |

Таблиці такого вигляду називаються таблицями спряженості.

Для розрахунку  $\chi^2$ -критерію Пірсона в програмі MS Excel, необхідно:

1) побудувати таблицю абсолютних частот значень змінної для двох вибірок;

2) побудувати таблицю теоретичних частот, обчисливши теоретичні частоти за формулою (2);

3) обчислити емпіричне значення критерію за формулою (1);

4) порівняти отримане емпіричне значення з критичними значеннями.

Критичні значення критерію можна знайти за допомогою таблиці критичних значень або обчислити за допомогою вбудованої функції ХИ2ОБР, підставивши в якості аргументів бажаний рівень значущості та число степенів свободи. Наприклад, функція =ХИ2ОБР(0,05;3) обчислить критичне значення  $\chi^2$ -критерію Пірсона на рівні значущості  $p \leq 0,05$  і числа степенів свободи  $k=3$ .

Здійснити аналіз даних за допомогою  $\chi^2$ -критерію Пірсона можна не розраховуючи емпіричного значення критерію, або не знаходячи критичних значень. Для цього використовуються функції ХИ2ТЕСТ та ХИ2РАСП відповідно.

Функція ХИ2ТЕСТ розраховує ймовірність однакового розподілу двох ознак, тобто ймовірність того, що у розподілі змінної в двох вибірках не буде значущих відмінностей (аналогічно функції ТТЕСТ). У якості аргументів необхідно зазначити фактичний (абсолютні емпіричні частоти) та очікуваний інтервал (абсолютні теоретичні частоти). Наприклад: =ХИ2ТЕСТ(B2:C4;F2:F4).

Функція ХИ2РАСП розраховує ймовірність однакового розподілу двох ознак для відомого емпіричного значення критерію та числа степенів свободи.

Наприклад: функція =ХИ2РАСП(В5;3) обчислить ймовірність, якщо емпіричне значення критерію міститься у комірці В5, а число степенів свободи дорівнює 3.

Розглянемо порядок виконання даного прикладу в програмі Excel.

1. Скопіюємо таблицю з вихідними емпіричними даними на аркуш Excel:

|   | A            | B       | C       | D    | E | F |
|---|--------------|---------|---------|------|---|---|
| 1 |              | 1 школа | 2 школа | Сума |   |   |
| 2 | Задоволені   | 15      | 7       | 22   |   |   |
| 3 | Незадоволені | 5       | 8       | 13   |   |   |
| 4 | Сума         | 20      | 15      | 35   |   |   |
| 5 |              |         |         |      |   |   |
| 6 |              |         |         |      |   |   |

У клітинках таблиці відображені абсолютні емпіричні частоти (кількість респондентів у двох школах, розподілених за параметром задоволеності психологічною службою).

2. Побудуємо таблицю теоретичних частот. Для цього використаємо формулу:

$$\text{Теоретична частота} = \frac{\text{сума відповідного рядка} \times \text{сума відповідного стовбця}}{\text{кількість респондентів у двох вибірках}}$$

Наприклад:

$$\begin{aligned} \text{Т. частота 1 рядка та 1 стовбця} &= \frac{\text{сума 1 рядка} \times \text{сума 1 стовбця}}{\text{кількість респондентів у двох вибірках}} = \\ &= \frac{22 \times 20}{35} \approx 12,57 \end{aligned}$$

$$\begin{aligned} \text{Т. частота 1 рядка та 2 стовбця} &= \frac{\text{сума 1 рядка} \times \text{сума 2 стовбця}}{\text{кількість респондентів у двох вибірках}} = \\ &= \frac{22 \times 15}{35} \approx 9,43 \end{aligned}$$

Для розрахунку теоретичних частот в Excel, у відповідні клітинки таблиці вводяться вищезазначені формули:

|              |              |                             |         |              |   |                    |              |         |         |   |
|--------------|--------------|-----------------------------|---------|--------------|---|--------------------|--------------|---------|---------|---|
| буфер обмена |              | Шрифт                       |         | Выравнивание |   |                    |              |         |         |   |
| H2           |              | $\times$ $\checkmark$ $f_x$ |         | $=D2*B4/D4$  |   |                    |              |         |         |   |
|              | A            | B                           | C       | D            | E | F                  | G            | H       | I       | J |
| 1            |              | 1 школа                     | 2 школа | Сума         |   | Теоретичні частоти |              | 1 школа | 2 школа |   |
| 2            | Задоволені   | 15                          | 7       | 22           |   |                    | Задоволені   | 12,57   | 9,43    |   |
| 3            | Незадоволені | 5                           | 8       | 13           |   |                    | Незадоволені | 7,43    | 5,57    |   |
| 4            | Сума         | 20                          | 15      | 35           |   |                    |              |         |         |   |
| 5            |              |                             |         |              |   |                    |              |         |         |   |
| 5            |              |                             |         |              |   |                    |              |         |         |   |

|              |              |                                                                                                                                                                                                                                                       |         |                                            |   |                    |              |         |         |   |   |
|--------------|--------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------|--------------------------------------------|---|--------------------|--------------|---------|---------|---|---|
| Буфер обмена |              | Шрифт                                                                                                                                                                                                                                                 |         | Выравнивание                               |   | Число              |              |         |         |   |   |
| I2           |              |    |         | <div><div></div><div>=D2*C4/D4</div></div> |   |                    |              |         |         |   |   |
|              | A            | B                                                                                                                                                                                                                                                     | C       | D                                          | E | F                  | G            | H       | I       | J | K |
| 1            |              | 1 школа                                                                                                                                                                                                                                               | 2 школа | Сума                                       |   | Теоретичні частоти |              | 1 школа | 2 школа |   |   |
| 2            | Задоволені   | 15                                                                                                                                                                                                                                                    | 7       | 22                                         |   |                    | Задоволені   | 12,57   | 9,43    |   |   |
| 3            | Незадоволені | 5                                                                                                                                                                                                                                                     | 8       | 13                                         |   |                    | Незадоволені | 7,43    | 5,57    |   |   |
| 4            | Сума         | 20                                                                                                                                                                                                                                                    | 15      | 35                                         |   |                    |              |         |         |   |   |
| 5            |              |                                                                                                                                                                                                                                                       |         |                                            |   |                    |              |         |         |   |   |
| 6            |              |                                                                                                                                                                                                                                                       |         |                                            |   |                    |              |         |         |   |   |

Буфер обмена

Шрифт

Выворивание

Числ

H3

✕

✓

*f<sub>x</sub>*

=D3\*B4/D4

|   | A            | B       | C       | D    | E | F                  | G            | H       | I       | J | K |
|---|--------------|---------|---------|------|---|--------------------|--------------|---------|---------|---|---|
| 1 |              | 1 школа | 2 школа | Сума |   | Теоретичні частоти |              | 1 школа | 2 школа |   |   |
| 2 | Задоволені   | 15      | 7       | 22   |   |                    | Задоволені   | 12,57   | 9,43    |   |   |
| 3 | Незадоволені | 5       | 8       | 13   |   |                    | Незадоволені | 7,43    | 5,57    |   |   |
| 4 | Сума         | 20      | 15      | 35   |   |                    |              |         |         |   |   |
| 5 |              |         |         |      |   |                    |              |         |         |   |   |
| 6 |              |         |         |      |   |                    |              |         |         |   |   |

|              |              |                             |         |              |   |                    |              |         |         |   |   |
|--------------|--------------|-----------------------------|---------|--------------|---|--------------------|--------------|---------|---------|---|---|
| Буфер обмена |              | Шрифт                       |         | Выравнивание |   | Число              |              |         |         |   |   |
| I3           |              | $\times$ $\checkmark$ $f_x$ |         | $=D3*C4/D4$  |   |                    |              |         |         |   |   |
|              | A            | B                           | C       | D            | E | F                  | G            | H       | I       | J | K |
| 1            |              | 1 школа                     | 2 школа | Сума         |   | Теоретичні частоти |              | 1 школа | 2 школа |   |   |
| 2            | Задоволені   | 15                          | 7       | 22           |   |                    | Задоволені   | 12,57   | 9,43    |   |   |
| 3            | Незадоволені | 5                           | 8       | 13           |   |                    | Незадоволені | 7,43    | 5,57    |   |   |
| 4            | Сума         | 20                          | 15      | 35           |   |                    |              |         |         |   |   |
| 5            |              |                             |         |              |   |                    |              |         |         |   |   |
| 6            |              |                             |         |              |   |                    |              |         |         |   |   |
| 7            |              |                             |         |              |   |                    |              |         |         |   |   |

3. У пусу клітинку введемо функцію ХИ2ТЕСТ (вставити формулу → статистичні). В якості аргументу «Фактичний інтервал» введемо діапазон даних вихідної таблиці (без сум). В якості аргументу «Очікуваний інтервал» введемо діапазон даних з таблиці теоретичних частот.

Библиотека функций | Определенные имен

G5 | X | ✓ | fx | =ХИ2.ТЕСТ(B2:C3;H2:I3)

|   | A            | B       | C       | D    | E | F                  | G            | H       | I       | J | K | L |
|---|--------------|---------|---------|------|---|--------------------|--------------|---------|---------|---|---|---|
| 1 |              | 1 школа | 2 школа | Сума |   | Теоретичні частоти |              | 1 школа | 2 школа |   |   |   |
| 2 | Задоволені   | 15      | 7       | 22   |   |                    | Задоволені   | 12,57   | 9,43    |   |   |   |
| 3 | Незадоволені | 5       | 8       | 13   |   |                    | Незадоволені | 7,43    | 5,57    |   |   |   |
| 4 | Сума         | 20      | 15      | 35   |   |                    |              |         |         |   |   |   |
| 5 |              |         |         |      |   |                    | C3;H2:I3     |         |         |   |   |   |

Аргументы функции

ХИ2.ТЕСТ

Фактический\_интервал B2:C3 = {15;7;5;8}

Ожидаемый\_интервал H2:I3 = {12,5714285714286;9,42857142857143;7,428571...}

= 0,086023236

Возвращает тест на независимость: значение распределения хи-квадрат для статистического распределения и соответствующего числа степеней свободы.

Фактический\_интервал диапазон, содержащий наблюдения, подлежащие сравнению с ожидаемыми значениями.

Значение: 0,086023236

[Справка по этой функции](#)

OK Отмена

Після натискання ОК отримаємо результат – ймовірність схожості двох вибірок. Пам'ятаємо, що гранична допустима ймовірність співпадає з рівнем значущості. Тобто для прийняття альтернативної гіпотези на рівні 0,01, ймовірність не повинна перевищувати 0,01. Для прийняття альтернативної гіпотези на рівні 0,05, ймовірність не повинна перевищувати 0,05. В нашому прикладі обчислена ймовірність становить 0,086, що більше 0,05. Тобто ми приймаємо нульову гіпотези, згідно якої дві досліджені школи не відрізняються між собою за рівнем задоволеності роботи психологічної служби.

**Завдання.** Виконайте даний приклад самостійно в програмі Excel, пройшовши всі вище перелічені етапи.

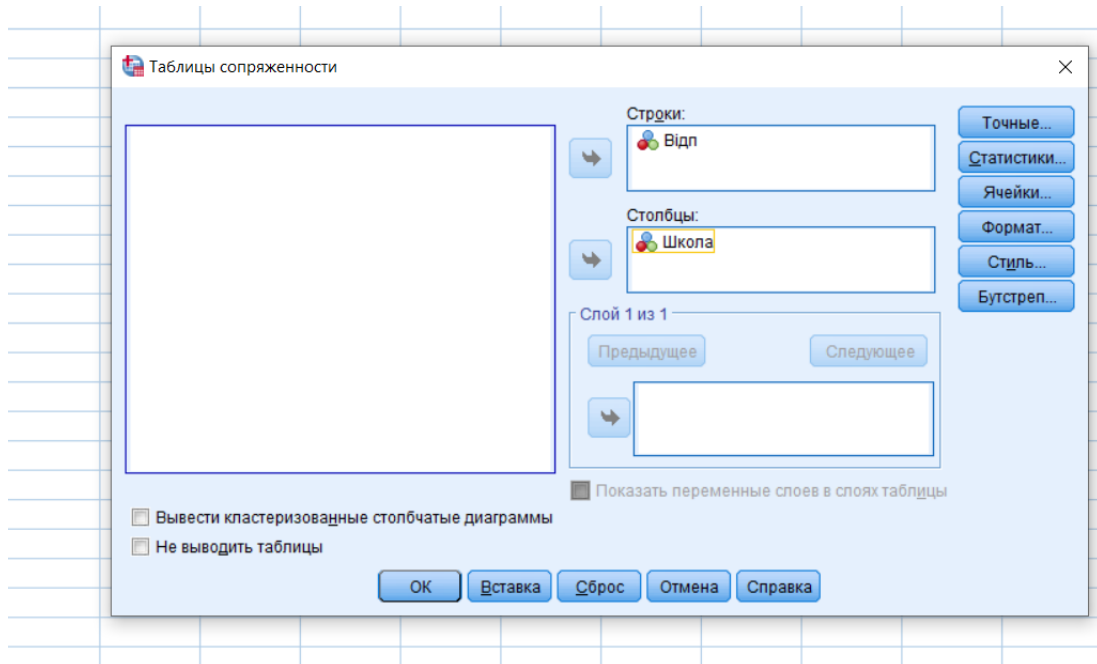
1. Для того, щоб обчислити хі-квадрат критерій у програмі SPSS, дані вихідної таблиці частот (в умові задачі) необхідно представити дещо в іншому форматі.

Відкриємо програму SPSS та заповнемо таблицю з двома стовбцями: «відповідь респондента» (1 – задоволений, 0 – незадоволений) та «школа» (1 або 2). Загальна вибірка з двох шкіл склала 35 респондентів (20 у першій школі та

15 у другій). У першій школі 15 задоволених респондентів, тож в перші 15 рядків вводимо цифру «1». У цій же школі 5 незадоволених респондентів, тож у наступні 5 рядків вводимо «0». Тепер у другий стовбець «школа» навпроти всіх введених цифр ставимо «1» (перша школа). Аналогічно заповнюємо таблицю для респондентів другої школи.

|    | Відп | Школа | пер. | пер. | пер. | п |
|----|------|-------|------|------|------|---|
| 1  | 1    | 1     |      |      |      |   |
| 2  | 1    | 1     |      |      |      |   |
| 3  | 1    | 1     |      |      |      |   |
| 4  | 1    | 1     |      |      |      |   |
| 5  | 1    | 1     |      |      |      |   |
| 6  | 1    | 1     |      |      |      |   |
| 7  | 1    | 1     |      |      |      |   |
| 8  | 1    | 1     |      |      |      |   |
| 9  | 1    | 1     |      |      |      |   |
| 10 | 1    | 1     |      |      |      |   |
| 11 | 1    | 1     |      |      |      |   |
| 12 | 1    | 1     |      |      |      |   |
| 13 | 1    | 1     |      |      |      |   |
| 14 | 1    | 1     |      |      |      |   |
| 15 | 1    | 1     |      |      |      |   |
| 16 | 0    | 1     |      |      |      |   |
| 17 | 0    | 1     |      |      |      |   |
| 18 | 0    | 1     |      |      |      |   |
| 19 | 0    | 1     |      |      |      |   |
| 20 | 0    | 1     |      |      |      |   |
| 21 | 1    | 2     |      |      |      |   |
| 22 | 1    | 2     |      |      |      |   |
| 23 | 1    | 2     |      |      |      |   |
| 24 | 1    | 2     |      |      |      |   |
| 25 | 1    | 2     |      |      |      |   |
| 26 | 1    | 2     |      |      |      |   |
| 27 | 1    | 2     |      |      |      |   |
| 28 | 0    | 2     |      |      |      |   |
| 29 | 0    | 2     |      |      |      |   |
| 30 | 0    | 2     |      |      |      |   |
| 31 | 0    | 2     |      |      |      |   |
| 32 | 0    | 2     |      |      |      |   |
| 33 | 0    | 2     |      |      |      |   |
| 34 | 0    | 2     |      |      |      |   |
| 35 | 0    | 2     |      |      |      |   |
| 36 |      |       |      |      |      |   |
| 37 |      |       |      |      |      |   |

2. Виберемо Аналіз – Описові статистики – Таблиці спряженості. У вікні, що відкрилось, змінну «Відповідь» введемо до поля «Строки», а змінну «Школа» - до поля «Столбцы».



3. Натиснемо кнопку «Статистики» та поставимо галочку навпроти параметру «Хи-квадрат». Натиснемо 2 рази ОК.

4. У вікні виводу отримаємо результати.

iPSS Statistics

| Преобразование Вставка Формат Анализ Прямой маркетинг Графика Сервис Окно Справка |          |          |             |          |       |          |
|-----------------------------------------------------------------------------------|----------|----------|-------------|----------|-------|----------|
|                                                                                   | Валидные |          | Пропущенные |          | Всего |          |
|                                                                                   | N        | Проценты | N           | Проценты | N     | Проценты |
| Відп * Школа                                                                      | 35       | 100,0%   | 0           | 0,0%     | 35    | 100,0%   |

Комбинационная таблица Відп \* Школа

|        |                      |  | Школа |      | Всего |
|--------|----------------------|--|-------|------|-------|
|        |                      |  | 1     | 2    |       |
| Відп 0 | Количество           |  | 5     | 8    | 13    |
|        | Ожидаемое количество |  | 7,4   | 5,6  | 13,0  |
| 1      | Количество           |  | 15    | 7    | 22    |
|        | Ожидаемое количество |  | 12,6  | 9,4  | 22,0  |
| Всего  | Количество           |  | 20    | 15   | 35    |
|        | Ожидаемое количество |  | 20,0  | 15,0 | 35,0  |

Частотна таблиця показує скільки задоволених та незадоволених в обох школах, та суми

Критерии хи-квадрат

|                                        | Значение           | ст. св. | Асимптотическая значимость (2-сторонняя) | Точная знач. (2-сторонняя) | Точная значимость (1-сторонняя) |
|----------------------------------------|--------------------|---------|------------------------------------------|----------------------------|---------------------------------|
| Хи-квадрат Пирсона                     | 2,947 <sup>a</sup> | 1       | ,086                                     |                            |                                 |
| Поправка на непрерывность <sup>b</sup> | 1,859              | 1       | ,173                                     |                            |                                 |
| Отношения правдоподобия                | 2,959              | 1       | ,085                                     |                            |                                 |
| Точный критерий Фишера                 |                    |         |                                          | ,157                       | ,087                            |
| Линейно-линейная связь                 | 2,863              | 1       | ,091                                     |                            |                                 |
| Количество допустимых наблюдений       | 35                 |         |                                          |                            |                                 |

a. Для числа ячеек 0 (0,0%) предполагается значение, меньше 5. Минимальное предполагаемое число равно 5,57.

b. Вычисляется только для таблицы 2x2

Емпіричне значення критерію

Ймовірність схожості вибірок (аналіз результату в Excel)

**Завдання.** Виконайте даний приклад самостійно в програмі SPSS, пройшовши всі вище перелічені етапи.

Таким чином:

Для застосування хі-критерію Пірсона необхідно виконувати наступні умови:

- 1) вимірювання може бути проведено в будь-якій шкалі;
- 2) об'єм вибірки не повинен бути меншим 30 респондентів;
- 3) теоретична частота для кожного інтервалу не менше 5;
- 4) вибірки повинні бути незалежними.

## Лекція № 7

### Тема. Регресійний аналіз

#### План

1. Поняття регресійного аналізу. Види регресійного аналізу.
2. Лінійна регресія. Коефіцієнти регресії. Рівняння регресії.
3. Інтерпретація результатів регресійного аналізу.
4. Здійснення регресійного аналізу у програмах MS Excel та SPSS.

#### **1. Поняття регресійного аналізу. Види регресійного аналізу.**

Регресійний аналіз – це статистичний метод дослідження впливу однієї або декількох незалежних змінних  $X$  на залежну змінну  $Y$ . При цьому незалежні змінні називаються регресорами або предикаторами, а залежні змінні – критеріальними. При цьому регресійний аналіз, так само як і кореляційний, не встановлює причинно-наслідковий зв'язок між змінними, а лише відображає ступінь зв'язку між ними. Регресійний аналіз не використовується для визначення наявності зв'язку між змінними, оскільки наявність такого зв'язку є посилкою для даного виду аналізу.

Регресійний аналіз відноситься до математичних методів прогнозування. За його допомогою можна спрогнозувати значення залежної змінної виходячи зі значень незалежної змінної.

#### **Приклад**

На деякій вибірці досліджували рівень інтелекту (IQ) батьків та їх дітей. Необхідно спрогнозувати рівень інтелекту дитини, знаючи рівень інтелекту батьків.

Спочатку за допомогою кореляційного аналізу виявляється взаємозв'язок між інтелектуальним рівнем батьків та дітей. За наявності значущого зв'язку можна говорити про застосування регресійного аналізу. У результаті регресійного аналізу виводиться формула, в яку можна підставити IQ батьків та отримати IQ дітей. Наприклад,  $IQ_{\text{дит}} = 0,667 \times IQ_{\text{род}} + 32,188$ ;  $IQ_{\text{дит}} = 0,667 \times 115 + 32,188 = 110$ .



У джерел регресійного аналізу стояв англійський вчений XIX ст. Френсіс Гальтон, який вважав, що антропологічні та інтелектуальні характеристики людей у наступних поколіннях вертаються до середніх для популяції показників. Так, якщо зріст батьків значно вище середнього, то скоріш за все зріст їх дітей також буде вище середнього, але нижче від зросту батьків. Отже, через декілька поколінь відбудеться регресія: значення зросту вернеться до середнього для популяції значення. Ті ж самі правила, на думку Гільтона, діють і по відношенню до інтелекту.

## **2. Лінійна регресія. Коефіцієнти регресії. Рівняння регресії.**

У найбільш простих випадках рівняння регресії має вигляд:  $Y=a_0+a_1X$ ;  $X=b_0+b_1Y$ .

Приклад.

Психологи виявили взаємозв'язок між успішністю навчання математиці  $Y$  та показником невербального інтелекту  $X$ . Було отримано наступне рівняння регресії:  $Y=1+0,025X$ . Припустимо, що показник невербального інтелекту учня дорівнює 132, тоді можна спрогнозувати середній показник успішності:  $Y=1+0,025*132=4,3$ . У іншого учня показник інтелекту 82, тоді:  $Y=1+0,025*82=3,05$ .

Регресійний аналіз в Excel.

Пакет аналіз, настройки. Коефіцієнт детермінації  $r^2$  від 0.5, 0.75.

Коефіцієнти:  $y$ -перетин –  $a_0$ , змінна  $x_1$  –  $a_1$

## Лекція № 8

### Тема. Факторний аналіз

#### План

1. Основні поняття факторного аналізу.
2. Методи обчислення факторного аналізу.

#### **1. Основні поняття факторного аналізу.**

В останні 30-40 років факторний аналіз набув значної популярності в психологічних і соціальних дослідженнях. Багато в чому цьому сприяла розробка Раймондом Кеттелла (Raymond B. Cattell) знаменитого 16-факторного особистісного опитувальника (16PF). Саме за допомогою факторного аналізу йому вдалося звести близько 4500 найменувань особистісних особливостей до 187 питань, які, в свою чергу, дозволяють виміряти 16 різних властивостей особистості.

Факторний аналіз дає можливість кількісно визначити щось безпосередньо не вимірюється, виходячи з декількох доступних виміру змінних.

Наприклад, характеристики «відвідує розважальні заходи», «багато розмовляє», «охоче йде на контакт з будь-яким незнайомою людиною» можуть служити оцінками якості «товариськість», яке безпосередньо не піддається кількісному вимірюванню. Факторний аналіз дозволяє встановити для великого числа вихідних ознак порівняно вузький набір «властивостей», що характеризують зв'язок між групами цих ознак і званих факторами. Процедура факторного аналізу складається з чотирьох основних стадій:

1. Обчислення кореляційної матриці для всіх змінних, що беруть участь в аналізі.
2. Витяг факторів.
3. Обертання факторів для створення спрощеної структури.
4. Інтерпретація факторів.

## 2. Методи обчислення факторного аналізу.

Обчислення кореляційної матриці. Перша операція, яка проводиться при виконанні факторного аналізу, — це обчислення кореляційної матриці для змінних, що беруть участь в аналізі. Вже за наявності такого початкового дії можна зробити висновок про те, що факторний аналіз заснований на взаємодії змінних. Для проведення факторного аналізу зовсім не обов'язково спеціально будувати кореляційну матрицю: при необхідності програма SPSS створить її сама на основі даних файлу. Іноді не потрібні навіть вихідні дані — достатньо мати кореляційну матрицю, яка в цьому випадку вводиться в командний файл SPSS. Однак подібний підхід дуже складний і в рамках цієї книги не розглядається.

При необхідності ви можете звернутися до керівництва користувача програми SPSS.

Витяг факторів. Витяг факторів є наступним етапом факторного аналізу. З математичної точки зору витяг чинників має певну аналогію з множинним регресійним аналізом. Першим кроком у множинному регресійному аналізі є вибір тієї незалежної змінної, яка обумовлює найбільшу частку дисперсії залежної змінної. Потім операція повторюється для незалежних змінних до тих пір, поки частка дисперсії не перестане бути значимою. У факторному аналізі існує аналогічна процедура.

Витяг фактора починається з підрахунку сумарного розкиду значень всіх що беруть участь в аналізі змінних (дана величина чимось схожа на загальну суму квадратів). Для цього «сумарного розкиду» непросто підібрати логічну інтерпретацію, однак він є цілком строго певною математичною величиною. Першої завданням факторного аналізу є вибір взаємодіючих змінних, чия взаємна кореляція обумовлює найбільшу частку загальної дисперсії. Ці змінні утворюють перший фактор.

Потім перший фактор виключається, і з решти безлічі змінних знову вибираються ті, чия взаємодія визначає найбільшу частку, що залишилася загальної дисперсії.

Ці змінні утворюють другий фактор. Процедура вилучення факторів продовжується до тих пір, поки не буде вичерпана вся загальна дисперсія змінних. За замовчуванням у процедурі факторного аналізу кожна змінна має одиничне значення спільності. Цей показник дорівнює частці дисперсії змінної, обумовленої сукупним впливом факторів. Спільність можна порівняти з множинним коефіцієнтом кореляції  $R$ , що приймають значення 0 у випадку, якщо фактори не впливають на змінну, значення 1 у випадку, якщо дисперсія змінної цілком визначається виділеними факторами. Перед початком вилучення факторів одиничне значення спільності встановлено за промовчанням для змінних, які беруть участь у факторному аналізі.

Після того, як процедура отримує перший фактор, навпроти його номера з'являється його власне значення, наприклад поруч з числом 1 з'являється значення 5,13312. Власне значення фактора пропорційно частці в загальній дисперсії, обумовленої даним фактором (зверніть увагу, що тут мова йде про чинник, а не про змінної, як було у випадку спільності). Перше власне значення завжди є найбільшим і перевищує одиницю; це визначається алгоритмом роботи процедури, згідно з якою фактори вилучаються у порядку зменшення їх впливу на дисперсію змінних. Потім обчислюється відсоток дисперсії, обумовлюваний даним фактором і рівний відношенню власного значення фактора до кількості змінних, а також відповідний кумулятивний (накопичений) відсоток. З отриманням кожного нового фактора власні значення зменшуються, а кумулятивний відсоток наближається до 100.

Зверніть увагу, що в контексті процедури вилучення факторів жодного разу не прозвучав термін «значимість»: на відміну від регресійного аналізу витяг факторів відбувається до тих пір, поки не буде вичерпана вся дисперсія змінних, незалежно від того, чи є значимим вплив фактора на дисперсію чи ні.

Вибір і обертання факторів. За дуже рідкісними винятками для дослідника не представляють інтересу все витягнуті фактори. Якщо чинників виявиться стільки ж, скільки вихідних змінних, факторний аналіз втрачає сенс, оскільки його метою є скорочення вихідного набору змінних.

Отже, потрібно прийняти рішення, які з факторів слід залишити для подальшого аналізу. Тут, в першу чергу, рекомендується керуватися здоровим глуздом і залишати ті фактори, які мають зрозумілу теоретичну або логічну інтерпретацію. Однак не завжди представляється можливим заздалегідь встановити призначення кожного фактора, і тому дослідники на першому етапі зазвичай використовують формальні критерії. За замовчуванням при виконанні команди Factor (Факторний аналіз) всі фактори, чії власні значення перевищують одиницю, зберігаються для подальшого аналізу.

Оскільки число факторів дорівнює числу змінних, лише для невеликої кількості факторів власні значення виявляються більше одиниці, означає виконання команди з параметрами за замовчуванням дозволяє радикально скоротити число факторів.

Існують і інші критерії виділення факторів (наприклад, критерій «кам'янистого осипу» Р. Кеттелла); крім того, ви можете вибрати чинники, ґрунтуючись на відомих вам особливостях конкретного файлу даних. У будь-якому випадку остаточне рішення про кількість чинників зазвичай приймається після інтерпретації факторів, отже, факторний аналіз передбачає неодноразове виділення різної кількості факторів. У розділі покрокових процедур розглянуті кілька варіантів виконання факторного аналізу, відмінні від прийнятого але за замовчуванням.

Наступним кроком після виділення факторів є їх обертання. Обертання потрібно тому, що спочатку структура факторів, будучи математично коректною, як правило, важка для інтерпретації.

Метою обертання є отримання простої структури, якій відповідає велике значення навантаження кожної змінної лише по одному чиннику і мале по всім іншим факторам. Навантаження відображає зв'язок між змінною і фактором,

будучи подобою коефіцієнта кореляції. Значення навантаження лежить в межах від -1 до 1. Ідеальна проста структура передбачає, що кожна змінна має нульові значення навантажень для всіх факторів, крім одного, для якого навантаження цієї неперемінної близька до 1 (-1).

Графічна інтерпретація процедури обертання представлена на рис. 20.1. До обертання (ліворуч) зірочки, відповідні змінним, розташовані на видаленні від осей факторів. Після повороту осей (праворуч) змінні виявляються поблизу осей, що відповідає максимальному навантаженню кожної змінної лише по одному чиннику. На практиці сувора орієнтація змінних вздовж осей факторів звичайно не досягається, проте операція повороту дозволяє наблизитися до бажаного результату. Після виконання обертання факторів програма SPSS містить висновки схожі графічні зображення, для зручності забезпечивши їх таблицею координат.

Обертання факторів анітрохи не впливає на математичну строгість аналізу: взаємне положення зірочок-змінних не змінюється при повороті осей. Спочатку процедура обертання виконувалася вручну, і дослідник самостійно ставив таке положення осей, при якому точки змінних максимально до них наближалися. Тепер SPSS дозволяє виконати кілька варіантів обертання, повертаючи осі так, щоб отримати просту структуру, яка задовольняє того чи іншого критерію.

Найбільш популярним варіантом обертання є метод Varimax. Варіант обертання Varimax є ортогональним, оскільки при такому обертанні осі зберігають своє взаємне розташування і од прямым кутом. Іноді можна отримати більш зручній просту структуру, якщо змінити кут між осями. Для цього призначені варіанти обертання Direct Oblimin і Promax. Наявність непрямого кута між осями чинників означає, що вони не є повністю незалежними один від одного.

Оскільки в реальних дослідженнях чинники дійсно не є абсолютно незалежними, відхилення кута між осями від прямого при обертанні цілком допустимо. Неортогональної обертання - досить складна процедура, і для її застосування необхідно чітко уявляти собі суть того, що відбувається. В цілому

аналогічну рекомендацію можна дати щодо факторного аналізу взагалі: не слід застосовувати його без попереднього вивчення відповідних розділів статистики. Варіанти обертання Direct Oblimin і Promax описані в розділі покрокових процедур, однак ми не демонструємо їх результати в розділі «Представлення результатів» зважаючи на складність останніх.

Інтерпретація факторів. Отже, нехай в деякій ситуації (близькою до ідеальної) шляхом обертання ми домоглися того, що значення навантаження для даного чинника є великим (більше 0,5), а для інших факторів - малим (менше 0,2); крім того, ми чітко уявляємо сенс нашого фактора, тобто те, що він вимірює.

Зрозуміло, в більшості досліджень змінні можуть взаємодіяти з «непотрібним» фактором, а нерідко таких факторів може бути кілька. Як правило, дослідник не обмежується тільки числовими результатами факторного аналізу; необхідною умовою успіху факторного аналізу є розуміння змістовної специфіки конкретних даних і взаємозв'язків між ними.

## Лекція № 9

### Тема. Кластерний аналіз

#### План

1. Поняття кластерного аналізу.
2. Методи здійснення та інтерпретації кластерного аналізу.

#### **1. Поняття кластерного аналізу.**

Основна перевага багатовимірного шкалювання — представлення великих масивів даних про відмінність об'єктів в наочному, доступному для інтерпретації графічному вигляді. При багатовимірному шкалюванні матриця відмінностей між об'єктами (обчислені, наприклад, але їх експертними оцінками) представляється у вигляді одно-, двох - або тривимірного графічного зображення взаємного розташування цих об'єктів. Хоча доступно і більше трьох вимірювань, ця можливість рідко застосовується на практиці.

Основною перевагою багатовимірного шкалювання є можливість дуже наочного візуального порівняння об'єктів аналізу. Якщо дві точки на зображенні віддалені один від одного, то між відповідними об'єктами є значна розбіжність; навпаки, близькість точок говорить про подібність об'єктів. Багатовимірне шкалювання має багато спільних рис з факторним аналізом. Так само як і при факторному аналізі, створюється система координат простору, в якому визначається розташування точок. Так само як і при факторному аналізі, відбувається зниження розмірності і спрощення даних. Однак при факторному аналізі зазвичай використовуються коефіцієнти кореляції, а при багатовимірному шкалюванні — заходи відмінності між об'єктами.

Нарешті, у факторному аналізі найбільший інтерес викликають кути між точками, що представляють дані, а в багатовимірному шкалюванні ключовий величиною є відстань між цими точками. Крім факторного аналізу багатовимірне шкалювання має кілька спільних рис з кластерним аналізом. В обох випадках



аналізується відстань між об'єктами; однак при кластерному аналізі типовою є кількісна процедура об'єднання об'єктів в групи (кластери), а при багатовимірному шкалюванні якісний аналіз об'єктів проводиться візуально за допомогою діаграми. Процедура багатовимірного шкалювання SPSS, що має історичну назву ALSCAL, фактично не є однією програмою, а являє собою набір невеликих процедур, кожна з яких відповідає своєму типу даних. У цій главі ми проведемо кілька аналізів для різних типів даних, намагаючись, але можливості, коротко і ясно викласти зміст супутніх термінів.

Якщо ви стикаєтеся з багатовимірним шкалюванням вперше, то рекомендуємо звернутися до додаткової літератури. У першому прикладі ми опрацюємо гіпотетичну соціограму для групи учнів. У цьому прикладі їх кількісні оцінки відносин один до одного будуть перетворені у графічне зображення взаємного розташування учнів. У другому прикладі ми розглянемо результати тестування учнів за п'ятьма показниками і представимо графічно відмінності між учнями на плоскому зображенні. Нарешті, третій приклад буде представляти собою невелике дослідження сприйняття та розуміння студентами п'яти багатовимірних методів статистичного аналізу. Найчастіше опис нового статистичного методу зручно проводити шляхом його порівняння з іншим методом. Саме таким чином ми і почнемо розгляд ієрархічного кластерного аналізу, порівнюючи його з факторним аналізом. При численних загальних рисах між зазначеними статистичними методами існує чимало відмінностей. Як звичайно, після теоретичної частини підуть приклади практичної реалізації статистичних методів засобами SPSS, оформлені у вигляді покрокових процедур.

Порівняння кластерного та факторного аналізів. Головна схожість між кластерним і факторним аналізами полягає в тому, що той і інший призначені для переходу від вихідної сукупності безлічі змінних (або об'єктів) до істотно меншому числу факторів (кластерів). Тим не менш реалізація статистичних процедур і інтерпретація результатів для двох типів аналізу розрізняються; п'ять основних відмінностей наведено далі. - Метою факторного аналізу є заміна великої кількості вихідних змінних меншим числом факторів. Кластерний аналіз,

як правило, застосовується для того, щоб зменшити кількість об'єктів шляхом їх групування. Іншими словами, у процедурі кластерного аналізу зазвичай змінні не групуються, а виступають в якості критеріїв для групування об'єктів. Так, в прикладі факторного аналізу 11 субтестів інтелекту (змінних) були зведені до трьох чинників, кожен з яких об'єднав кілька споріднених вихідних змінних. Кластерний аналіз застосовується зазвичай для виділення груп об'єктів, виходячи з їх подібності за вимірними ознаками. Стосовно до прикладу з 11 субтестами інтелекту типовою задачею кластерного аналізу була б класифікація учнів (об'єктів) таким чином, щоб але вимірним 11 показниками всередині кожної групи об'єкти були б схожі один на одного, ніж на об'єкти з інших груп. Групи об'єктів, виділені в результаті кластерного аналізу на основі заданої міри подібності між об'єктами, що називаються кластерами.

## **2. Методи здійснення та інтерпретації кластерного аналізу.**

Заявлені в попередньому пункті відмінності між кластерним і факторним варіантами аналізу з усією повнотою категоричності можуть бути віднесені лише до попередніх версій SPSS.

Починаючи з версії SPSS 10.0, програма дозволяє з однаковим успіхом проводити кластерний аналіз не лише об'єктів, а й змінних. В останньому випадку кластерний аналіз може виступати як більш простий і нерідко більш ефективний аналог факторного аналізу. У розділі покрокових процедур ми продемонструємо обидва варіанти кластерного аналізу. Дії, що виконуються в ході статистичних операцій в кожному з варіантів аналізу, принципово різняться. У факторному аналізі на кожному етапі вилучення фактора для кожної змінної підраховується частка дисперсії, яка обумовлена впливом цього чинника. При кластерному аналізі обчислюється відстань між поточним об'єктом і всіма іншими об'єктами, і утворює кластер та пара, для якої відстань виявилось найменшим. Подібним чином кожен об'єкт або групується з іншим об'єктом, або включається до складу існуючого кластера. Процес кластеризації кінцевим і

триває до тих пір, поки всі об'єкти не будуть об'єднані в один кластер. Зрозуміло, подібний результат в загальному випадку не має сенсу, і дослідник повинен самостійно визначити, в який момент кластеризація повинна бути припинена. У контексті кластерного аналізу особливе місце займає один із його видів, званий ієрархічним кластерним аналізом. В SPSS він реалізується з допомогою команди Hierarchical Cluster (Ієрархічна кластеризація).

Цей вид кластерного аналізу частіше використовується в економіці, соціології, політології, ніж у психології. Психологи зазвичай аналізують змінні з метою знайти статистичні зв'язки між ними; ці зв'язки, як правило, вказують на схожість між тими чи іншими досліджуваними факторами. Розподіл вибірки па групи в психологічних аналізах рідко представляє інтерес; у випадках, коли це виявляється необхідним, психологи віддають перевагу дискримінантному, а не кластерним аналізом. Оскільки кластеризація змінних виявляється вельми доступною операцією, було б цікаво порівняти її результати з результатами більш складного факторного аналізу. Як і у випадку факторного аналізу, виконання кластерного аналізу та його результати залежать від ряду параметрів: способу обчислення відстані між об'єктами, кластеризації індивідуальних об'єктів і т. д. Етапи кластерного аналізу.

Кластерний аналіз виконується за кілька етапів, що призводять до кінцевого результату. Спочатку ми розглянемо приклад, створений спеціально, щоб показати суть кластерного аналізу. Зазначимо, що кластерний аналіз непридатний до файлів даних, що використовувалися раніше, оскільки при їх складанні основну увагу було приділено змісту і зв'язків між змінними, а вміст об'єктів (тобто інформація, що стосується суб'єктів) практично не грало ролі. Для демонстрації кластерного аналізу нами був підготовлений файл cars.sav, що містить гіпотетичні дані про 15 автомобілях різних марок, виставлених на продаж. Файл має структуру, відповідну для наочної ілюстрації кластерного аналізу. Отже, виділяють кілька етапів кластерного аналізу. 1. Вибір показників-критеріїв для кластеризації.

У нашому прикладі кластеризація буде здійснюватися за наступними змінними: ціна (вартість),  $t\_cost$  (експертна оцінка технічного стану за 10-бальною шкалою), вік (кількість років експлуатації), пробіг (пройдений кілометраж з початку експлуатації). 2. Вибір способу вимірювання відстані між об'єктами, або кластерами (спочатку вважається, що кожен об'єкт відповідає одному кластеру). За замовчуванням використовується квадрат Евклідової відстані, згідно з яким відстань між об'єктами дорівнює сумі квадратів різниць між значеннями однойменних змінних об'єктів. Припустимо, що марка автомобіля А має показники технічного стану і віку 5 і 6, а марка — відповідно 7 і 4. У цьому випадку відстань між марками обчислюється наступним чином:  $(5 - 7)^2 + (6 - 4)^2 = 8$ . При виконанні аналізу сума квадратів різниць обчислюється для всіх змінних. Одержувані відстані використовуються програмою при формуванні кластерів.

Крім Евклідова існують і інші види відстаней, обчислювані за іншими формулами, проте ми не будемо на них зупинятися. При необхідності зверніться до керівництва користувача SPSS. Щодо обчислення відстані може виникнути наступне запитання: чи буде адекватним результат кластерного аналізу у тому випадку, якщо змінні мають різні шкали вимірювання? Так, всі змінні файлу `car.sav` мають різні шкали. Для вирішення проблеми шкалювання в SPSS використовується стандартизація, зокрема, її простий метод — нормалізація змінних, що приводить всі змінні до стандартної z-шкалою (середнє дорівнює 0, стандартне відхилення — 1). Крім однаковою шкали нормалізовані змінні також мають рівні ваги. У разі якщо всі вихідні дані мають одну і ту ж шкалу вимірювання або ваги змінних за змістом повинні бути різними, стандартизацію змінних проводити не потрібно.

Формування кластерів. Існує два основних методи формування кластерів: метод злиття і метод дроблення. В першому випадку вихідні кластери збільшуються шляхом об'єднання до тих пір, поки не буде сформований єдиний кластер, що містить всі дані. Метод дроблення заснований на зворотної операції: спочатку всі дані об'єднуються в один кластер, який потім ділиться на частини

до тих пір, поки не буде досягнутий бажаний результат. За замовчуванням програмою SPSS використовується метод злиття. У методі злиття передбачено кілька способів об'єднання об'єктів.

Спосіб, використовуваний за умовчанням, називається міжгрупових зв'язуванням, або зв'язуванням середніх всередині груп. SPSS обчислює найменше середнє значення відстані між усіма нарами груп і об'єднує дві групи, які опинилися найбільш близькими. На першому кроці, коли всі кластери являють собою поодинокі об'єкти, дана операція зводиться до звичайного попарному порівнянні відстаней між об'єктами. Термін «середнє значення» набуває сенсу лише на другому етапі, коли сформовані кластери, що містять більше одного об'єкта. Так, в орємо прикладі на початковому етапі є 15 кластерів (об'єктів); спочатку в кластер об'єднуються два об'єкти з найменшою відстанню один від одного. Потім підрахунок відстаней повторюється, і в кластер об'єднується ще одна пара змінних.

На другому етапі ви отримаєте або 13 вільних об'єктів і 1 кластер, який об'єднує 2 об'єкти, або 11 вільних об'єктів і 2 кластера по 2 об'єкти в кожному. В кінцевому рахунку всі об'єкти опиняться в одному великому кластері. Існують і інші методи об'єднання об'єктів, проте ми не станемо розглядати їх в цій книзі. При необхідності зверніться до керівництва користувача SPSS.

Інтерпретація результатів. Як і в випадку факторного аналізу, бажане число кластерів і оцінка результатів аналізу залежать від цілей дослідника. Для розглянутого прикладу нам представляється найкращим число кластерів, рівну 3. Як показує аналіз, все марки можна розділити на 3 групи: перша група має високу вартість (середнє значення - 15 230), недовгий термін експлуатації (4 роки) і середній пробіг (85 400 км). Друга група має середню вартість, невеликий пробіг, максимальний вік, але гарний технічний стан. Третя група містить недорогі моделі з великим пробігом і невисоким рейтингом технічного стану.

## Список рекомендованої літератури

### Основна:

1. Ермолаев О.Ю. Математическая статистика для психологов : Учебник. 2-е изд., испр. Москва : Московский психолого-социальный институт: Флинта, 2003. 336 с.
2. Климчук В.О. Математичні методи у психології. Навчальний посібник для студентів психологічних спеціальностей. Київ : Освіта України. 2009. 288 с.
3. Літнарів Р.М. Основи математичної статистики у психології : Навчальний посібник. Ч.3. Рівне : МЕРУ, 2006. 49 с.
4. Руденко В.М., Руденко Н.М. Математичні методи в психології : підручник. Київ : Академвидав, 2009. 384 с.
5. Татьянчиков А.О. Методичні рекомендації до виконання лабораторних робіт з курсу «Методи психологічного дослідження: математичні методи в психології». Одеса : Вид-во Університету Ушинського, 2019. 38 с.
6. Телейко А.Б. Чорней Р.К. Математико-статистичні методи в соціології та психології : Навч. посібник. Київ : МАУП, 2007. 424 с.

### Допоміжна:

1. Боснюк В.Ф. Математичні методи в психології курс лекцій. – Харків 2016.
2. Наследов А.Д. Математические методы психологического исследования. Анализ и интерпретация данных : Учебное пособие. Санкт-Петербург : Речь, 2004. 392 с.
3. Резник А.Д. Книга для тех, кто не любит статистику, но вынужден ею пользоваться. Непараметрическая статистика в примерах, упражнениях и рисунках. Санкт-Петербург : Речь, 2008. 265 с. Сидоренко Е.В. Методы

математической обработки в психологии. Санкт-Петербург : Речь, 2002. 350 с.

4. Татьянчиков А.О. Формування розумових операцій як засіб адаптації підлітків до навчання в основній школі : дис. ... канд. психол. наук : 19.00.07. Слов'янськ, 2013. 252 с.
5. Татьянчиков А.О. Динаміка факторів адаптації підлітків на етапі переходу до навчання в основній школі. Наука і освіта : науково-практичний журнал. Одеса, 2016. № 7/CXXXXVIII. С. 20-25.
6. Татьянчиков А.О. Динаміка шкільної тривожності учнів на етапі переходу до навчання в основній школі. Вісник Харківського національного університету імені В. Н. Каразіна. Харків, 2013. Вип. 1065. С. 156-159.
7. Foster, G., Lane D.; Scott D., Hebl M. and other. An Introduction to Psychological Statistics. University of Missouri, St. Louis. 2018. 271 p.