

Міністерство освіти і науки України
Дніпропетровський національний університет

В.Л. Борщ, Ю.П. Совіт

**МАТЕМАТИЧНІ МЕТОДИ В БІОЛОГІЇ.
ОСНОВИ ТЕОРІЇ ЙМОВІРНОСТЕЙ**

Конспект лекцій

Міністерство освіти і науки України
Дніпропетровський національний університет

В.Л. Борщ, Ю.П. Совіт

**МАТЕМАТИЧНІ МЕТОДИ В БІОЛОГІЇ.
ОСНОВИ ТЕОРІЇ ЙМОВІРНОСТЕЙ**

Конспект лекцій

Дніпропетровськ
РВВ ДНУ
2007

УДК 519.22/.25 (075.8)

Б83

Рецензенти: д-р фіз.-мат. наук, проф. М.В. Поляков
д-р фіз.-мат. наук, проф. С.Б. Вакарчук

Видання містить лекції з курсу “Математичні методи в біології”. Вони узагальнюють досвід викладання цієї дисципліни студентам біолого-екологічного факультету ДНУ. Означеним матеріалом можна послуговуватися в процесі вивчення інших дисциплін, де застосовується числова обробка даних дослідів.

Б83

Борщ В.Л. Математичні методи в біології. Основи теорії ймовірностей: Конспект лекцій / В.Л.Борщ, Ю.П.Совіт. – Д.: РВВ ДНУ, 2007. – 140 с.

Передмова

Пропоноване видання втілює узагальнений багаторічний досвід авторів у викладанні дисципліни „Математичні методи“ студентам біолого-екологічного факультету Дніпропетровського національного університету. За цей тривалий час змінювалися навчальні програми, кількість годин, перелік і зміст тем у складі дисципліни та ін. Не змінювалося лише ставлення авторів до викладання дисципліни, точніше до питання, яким чином, на якому підґрунті її викладати. Поза усяких сумнівів, можливими (і, до речі, опрацьованими) є такі варіанти викладання: рецептурний, що складається з інструкцій щодо застосування тих чи інших методів обробки даних без усяких доведень і навіть обґрунтувань, та подібний тому, що викладається математиками для математиків – аксіоматичний, хоча і суттєво скорочений. В основі останнього – попереднє введення базових первинних елементів (які не визначаються, а тільки називаються), а властивості первинних елементів і їх співвідношення задаються аксіомами. Стосовно аксіом завжди виникають проблеми повноти, надлишковості та відповідності фізичній реальності, але ці питання не так часто ставляться і на них не дуже часто відповідають. Дані питання торкаються умовчань стосовно вибору тієї чи іншої системи аксіом, що безпосередньо приводить до „науки, яка передує науці, що розглядається“, а це руйнує очевидні переваги аксіоматичного методу – можливість достатньо стислого і логічно послідовного викладення матеріалу. Аксіоматичний метод є класичним дедуктивним методом, тобто таким підходом, у якому часткові результати є наслідком загальних положень. Звернемо увагу на те, що цей метод є достатньо розповсюдженим, є навіть спроби застосування його до біології (Медников, 1987). Стосовно теорії ймовірностей достатньо послідовним критиком аксіоматичного підходу був Мізес (1930). Погляди Мізеса послідовно проаналізував Хінчин (1929, 1934), і цей аналіз ще й досі викликає зацікавленість. Ми обмежимося лише тим, що підкреслимо головні погляди Мізеса, який вважав, що ймовірність, основне поняття теорії ймовірностей, не може бути невизначеним, воно повинно бути визначеним і вимірювальним у досліді. Мізес також вважав, що відносна частота, за умови її дослідної стабілізації (так звана статистична ймовірність), і є фізичним визначенням ймовірності. Додамо, що дискусії на цю тему тривають і за наших часів. Наприклад, під час дискусії, що відбулася на сторінках журналу „Успехи физических наук“ (Кравцов, 1989; Алімов, 1992; Тутубалін, 1992), висловлювалася і обґрунтовувалася думка, що ймовірність – це все ж таки фізична величина. Саме на ідеях Мізеса щодо фізичності ми і засновуємо викладення матеріалу дисципліни „Математичні методи“ в наших лекціях. Цей підхід є цілком індуктивним. Індуктивний підхід ґрунтується на узагальненні фактів у вигляді теорії, тобто на русі від „часткового“ до „загального“.

Логічна й тематична побудова лекцій є така.

На початку подаються необхідні для подальшого викладення відомості та поняття з теорії множин і комбінаторики (теми 1 і 2). Для того, щоб вони дійсно стали зрозумілими, ми дали їх обґрунтування й доведення за допомогою великої кількості схем і діаграм. Тобто ми намагалися зробити ці поняття й відомості наочними, а ситуації, коли слід їх використовувати, – модельними.

Далі вводяться основні поняття і результати теорії ймовірностей стосовно випадкових подій (тема 3). Одним із ключових понять у цій темі є поняття досліду, подання та подальше використання якого здійснюється під впливом поглядів, викладених у підручнику „Теория вероятностей“ (Вентцель, 1970). Автори цього підручника є відомими фахівцями застосування ймовірнісних та статистичних методів у деяких технічних галузях, що й відрізняє викладення окремих розділів у підручнику від викладення для „чистих“ математиків. Звертаємо увагу на те, що в цій темі майже завжди при доведенні застосовується статистична ймовірність, і для наочності ми пропонуємо використовувати „лабораторний журнал“. Для тих студентів, у яких добре розвинене образне мислення, доведення дублюються із застосуванням геометричної ймовірності¹ і діаграм Венна (паралельне викладення, яке застосовується в більшості тем). Крім того, особлива увага приділяється такому важливому поняттю, як незалежність випадкових подій. Ми пропонуємо своє „визначення“ незалежності подій (взагалі явищ), яке, на наш погляд, можна застосовувати в досліді. На ньому ґрунтується будується визначення незалежності випадкових подій. У результаті одержуємо правило множення ймовірностей, яке впливає з визначення незалежності подій, а не є первинним визначенням.

У темі 4 мова йде про випадкові величини, даються пояснення їх відмінності від випадкових подій. Викладення цієї теми поєднане з викладенням вибіркового методу, тобто ми спираємося на фізичний характер ймовірності, намагаючись розвивати ідеї Мізеса. Особливу увагу слід звернути на те, що основні поняття в описі випадкових величин, а саме функція густини і функція розподілу, не вводяться, а “конструюються” за осмисленням результатів обробки первинних вибірових даних. Основним робочим поняттям тут є інтервальна випадкова величина, яка, на відміну від ідеальної моделі у вигляді неперервної випадкової величини, саме й спостерігається в досліді. Звичайно, маючи тільки обмежену кількість дослідних даних, а в уявному необмежено тривалому досліді тільки їх зліченню послідовність (відверто кажучи, ми не приховували, що реально їх кількість все ж таки скінченна, через обмеженість реальних генеральних сукупностей), ми не можемо абсолютно строго здійснювати граничні переходи від інтервальних величин до неперервних, тобто заповнювати „без проміжків“ скінченні інтервали. Отже, не всі граничні переходи вигляду $n \rightarrow \infty$, що вважаються здійсненими при диференціюванні й інтегруванні, можна вважати строгими. Однак у фізиці такий підхід приводить до правильних результатів. Наприклад, механіка суцільного середовища вдало поєднує ідею неперервності й математичний апарат диференціювання та інтегрування стосовно фізичного середовища, яке, по суті, є дискретним і зліченим. У пропонованому конспекті лекцій завдяки такому підходу впровадження багатьох понять стає значно простішим. Зокрема, це стосується парних понять, таких як статистики і параметри. Перші визначаються в досліді, тобто за вибірками, а другі – за вибірками максимального обсягу (така вибірка насправді одна), тобто за генеральною сукупністю.

¹ Цікаві для біологів приклади щодо геометричної ймовірності наводить Ю.І. Гільдерман (1991).

У процесі викладення цієї теми значна увага приділяється поняттю квантилів. Розв'язуючи за їх допомогою „пряму“ і „обернену“ задачі, ми наголошуємо на необхідності прийняття рішень через неоднозначність розв'язку „оберненої“ задачі. Це, на нашу думку, сприятиме виробленню в студентів уміння приймати рішення при перевірці статистичних гіпотез, особливо в тих випадках, коли треба вибирати між однобічними та двобічними критеріями перевірки та задавати рівень значущості. Можливо, саме через поширеність рецептурного стилю викладення методів обробки даних робота з квантилями супроводжується нерозумінням. Ми зробили спробу це виправити.

У темах 5 і 6 дані стандартні розподіли для дискретних і неперервних величин, тут майже відсутні відмінності від звичайних способів викладення.

Ми не ставили за мету навіть стосовно важливих результатів вдаватися до класичних доведень, натомість обмежувалися достатнім, на наш погляд, „здоровим глуздом“. У тих нечисленних випадках, де доведення зустрічаються, ми позначали їх завершення знаком ■. Знак ● застосовувався для позначки завершення прикладів, а знак ▲ – завершення зауважень. На зауваження, подані у конспекті лекцій, доцільно звернути увагу з такого приводу. Вони не є штучними, тобто їх не слід вважати навмисними виділеннями в тексті. Переважна їх більшість є відповідями на запитання студентів під час лекцій або аналізом їх помилок на практичних заняттях або в відповідях під час іспитів.

1. Елементи теорії множин

1.1. Поняття множини

Множина – будь-яка чітко визначена сукупність (набір, клас, група, об'єднання тощо) об'єктів довільної природи, що мають яку-небудь спільну ознаку. У визначенні акцент робиться на спільності об'єктів множини за деякою ознакою. Зрозуміло, що включити об'єкти з такою ознакою в дану множину можна тільки шляхом вибору їх із "більш загальної" множини, яка містить об'єкти і з іншими ознаками. Ця множина має особливу назву – **універсальна множина**. Надалі об'єкти будь-якої множини будемо називати **елементами множини**.

Отже, побудову будь-якої множини можна здійснити, якщо задані:

- 1) універсальна множина;
- 2) правило, що дозволяє відносно будь-якого елемента універсальної множини однозначно вирішити питання про належність його даній множині.

Приклад 1.1. Універсальна множина – студенти Дніпропетровського національного університету. З універсальної множини можна, наприклад, виділити множину студентів біолого-екологічного факультету. ●

Множини позначаються великими літерами латинського алфавіту ($A, B, C, \dots$), а елементи множин – малими літерами ($a, b, c, \dots$). Універсальна множина позначається літерою U . Для позначення того, що елемент a універсальної множини U є елементом множини A (ще говорять, що об'єкт a належить множині A) використовується запис $a \in A$. Запис $a \notin A$ означає, що елемент a універсальної множини U не належить множині A .

Якщо відома деяка властивість T , що дозволяє з елементів універсальної множини побудувати множину A , то множина A може бути описана в такий спосіб:

$$A = \{a \in U: T(a) - \text{істинне твердження}\}.$$

Цей запис читається так: множина A складається з тих елементів універсальної множини, для яких властивість T є істинна.

Приклад 1.2. Нехай універсальна множина U – всі люди, які живуть на Землі; множина A складається з людей, які мають певне захворювання. Оскільки захворювання виявляється за результатами діагностичного обстеження або проведення спеціального тесту (у даному прикладі це й буде властивість T), то множина може бути описана так, як зазначено вище. Важливо відзначити, що повний список елементів множини A може бути й не відомий (через надто велику кількість елементів універсальної множини), але відносно будь-якого елемента універсальної множини U (тобто будь-якої людини) однозначно вирішується питання про включення в A шляхом перевірки для нього виконання властивості T (людина або хвора, або здорова, точніше, тест на захворювання або позитивний, або негативний). ●

Іноді вдається перелічити всі елементи множини A , і це також є опис множини:

$$A = \{a_1, a_2, \dots, a_i, \dots, a_n\}.$$

Порожньою множиною називається множина, що не містить жодного елемента; вона позначається знаком $\emptyset$. Порожня множина може бути описана на елементах універсальної множини U за допомогою деякої властивості F , що не виконується для елементів U :

$$\emptyset = \{a \in U: F(a) - \text{хибне твердження}\}.$$

Приклад 1.3. Такі множини людей є порожніми:

$$\emptyset = \{a \in U: \text{зріст}(a) > 3 \text{ м}\};$$

$$\emptyset = \{a \in U: \text{вага}(a) > 500 \text{ кг}\};$$

причому U може бути обрана зовсім довільно, скажімо, як у прикладах 1.1 і 1.2.

За кількістю елементів множини поділяються на **скінченні** та **нескінченні**. У деяких випадках скінченні множини зручно вважати **практично нескінченними**, якщо часові, фізичні, фінансові або які-небудь інші причини не дозволяють перебрати всі елементи множини.

Приклад 1.2 (продовження). Універсальна множина U , поза всяким сумнівом, є скінченна (відома верхня оцінка чисельності населення Землі), але з погляду можливості проведення масового обстеження, опитування, перепису і т.ін. всіх елементів U її варто визнати практично нескінченною. ●

Зауваження 1.1. Поділ множин на скінченні та нескінченні не означає, що питання про визначення деякої множини зводиться до підрахунку кількості її елементів. Кількість елементів множини може бути не відома, але множина завжди має чітке визначення. Для множини A в прикладі 1.2 не можна вказати кількість елементів з ряду причин: універсальна множина є практично нескінченна, склад самої множини змінюється внаслідок природного приросту й зменшення населення й т.ін. Тобто чисельність множини взагалі не має ніякого відношення до визначення множини. ▲

1.2. Порівняння множин

Дві множини A і B є рівні, або збігаються, якщо вони містять ті ж елементи або, що те саме, складаються з тих же елементів. Усі порожні множини за визначенням вважаються рівними.

Приклад 1.3 (продовження). Обидві порожні множини в прикладі 1.3 є рівні, хоча вони визначені за допомогою **різних** (!) умов F (див. розділ 1.1). ●

Збіг множин позначається знаком рівності:

$$A = B.$$

Множина B називається **підмножиною** (частиною) множини A , якщо всі елементи B належать також і A , що позначається в такий спосіб:

$$B \subseteq A, \text{ або } A \supseteq B.$$

Очевидно, що кожна множина A є підмножиною універсальної множини U і самої себе: $A \subseteq U$, $A \subseteq A$. Вважається, що порожня множина є підмножиною будь-якої множини: $\emptyset \subseteq A$.

Множина B є **істинною (властивою) підмножиною** множини A , якщо: $B \subseteq A$, $B \neq \emptyset$, $B \neq A$. Те, що множина B є властивою підмножиною множини A , позначається так:

$$B \subset A, \text{ або } A \supset B.$$

Зауваження 1.2. Можна порівняти знаки $\subseteq$ і $\subset$, що використовуються для запису відношень між множинами, із знаками $\leq$ і $<$, які застосовуються для запису відношень між числами. ▲

Зауваження 1.3. Відношення між множинами встановлюються на основі властивостей, яким задовольняють їх елементи, а не за кількістю елементів. Інакше кажучи, із того, що кількість елементів множини B менша кількості елементів множини A , не випливає, що множина B є підмножиною множини A . ▲

Множини зручно зображувати за допомогою фігур на площині (їх називають **діаграмами Ейлера – Венна** або **діаграмами Венна**), точки яких зображують елементи множин. Універсальна множина зображується на діаграмі Венна більшою фігурою (наприклад, прямокутником), усередині якої розташовуються менші фігури – множини. Графічна інтерпретація відношень між множинами (рис. 1) часто виявляється зручною навіть тоді, коли самі множини (наприклад, множина кроликів) не мають нічого спільного з геометричними об'єктами (точками фігур, які зображують множини).

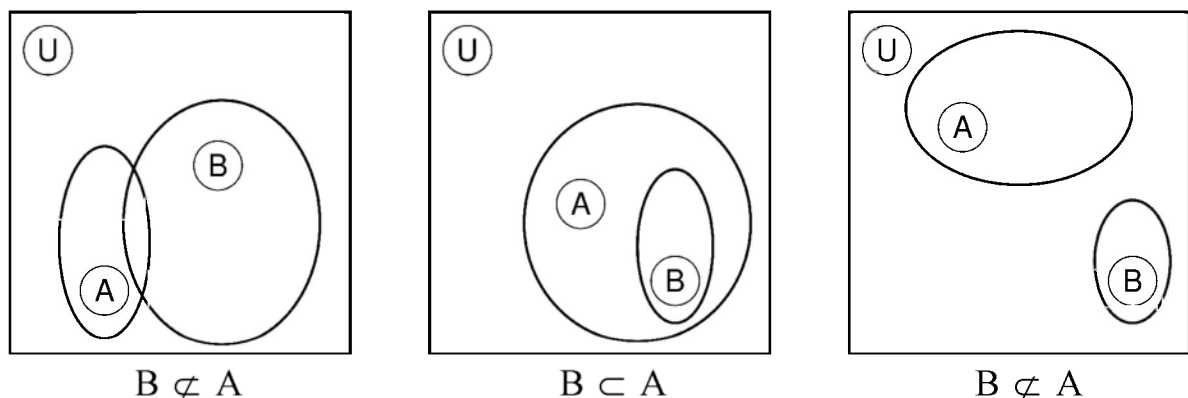


Рис. 1.1. Діаграми Венна

1.3. Дії над множинами

Сумою (об'єднанням) множин A і B називається множина

$$C = A + B,$$

що складається з усіх елементів U , що належать **або** A , **або** B , тобто хоча б одній із цих множин (у тому числі й тих, які належать **і** A , **і** B). Іншими словами, сума двох множин – це сукупність елементів, що належать хоча б одній із них. Сума є результатом дії, що називається додаванням, і позначається знаком “+”. Додавання для множин A і B , показаних на рис. 1.1, пояснюється на рис. 1.2 (множина C позначена сірим кольором).

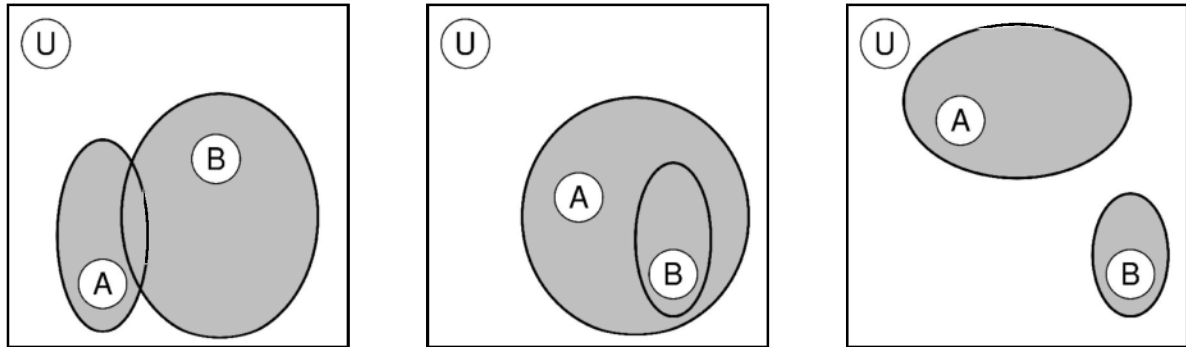


Рис. 1.2. Додавання множин

Приклад 1.4. Нехай у деякій популяції (універсальна множина U) визначені множини: $A = \{\text{люди, молодші 10 років}\}$, $B = \{\text{люди, старші 10 років і молодші 20 років}\}$, $C = \{\text{люди, старші 20 років}\}$, $D = \{\text{люди, старші 15 і молодші 25 років}\}$. Тоді $A+B = \{\text{люди, молодші 20 років}\}$, $B+C = \{\text{люди, старші 10 років}\}$, $A+C = \{\text{люди, молодші 10 або старші 20 років}\}$. ●

Аналогічно сумі двох множин визначається **сума (об'єднання) скінченної кількості множин**

$$C = \sum_{i=1}^n A_i = A_1 + A_2 + \dots + A_n$$

як множина всіх елементів, що належать хоча б одній із множин $A_i, i = \overline{1, n}$.

Добутком (перетином) множин A і B називається множина

$$D = A \cdot B = AB,$$

що складається з елементів U , які належать **і** множині A , **і** множині B , тобто спільних для A і B . Добуток є результат дії, що називається множенням і позначається знаком “·”, який, проте, використовується досить рідко. Множення для множин A і B , показаних на рис.1.1, пояснюється на рис.1.3 (множина D позначена сірим кольором).

Приклад 1.4 (продовження). $AB = \emptyset$, $BC = \emptyset$, $BD = \{\text{люди, старші 15 і молодші 20 років}\}$. ●

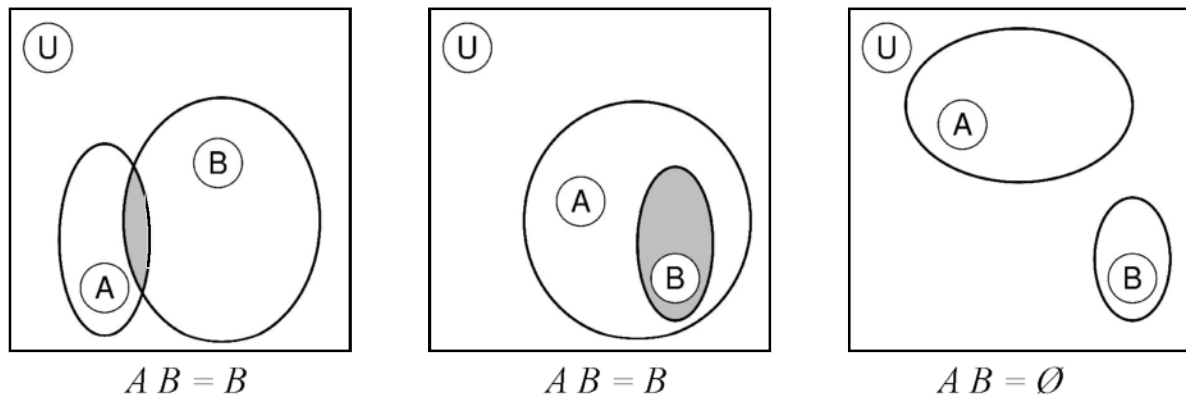


Рис. 1.3. Множення множин

Аналогічно добутку двох множин визначається **добуток (перетин) декількох множин**

$$D = \prod_{i=1}^n A_i = A_1 A_2 \dots A_n$$

як множина елементів, спільних для множин $A_i, i = \overline{1, n}$.

Зауваження 1.4. Легко запам'ятати, чим відрізняються сума й добуток множин, якщо звернути увагу на сполучники в їх визначеннях: **або** – для суми, **і** – для добутку. ▲

Зауваження 1.5. Поширена при вивченні суми й добутку множин помилка полягає в тому, що в їх визначеннях слово "елемент" "легко" замінюється словом "множина". Наприклад (це помилкове визначення!), сумою множин A і B називається множина C , що складається із множини A або множини B . Помилковість такого "визначення" полягає в невизначеності результату. Дійсно, із визначення множини (див. розд. 1.1) випливає, що можна (коли множина сформована) і необхідно (коли множина формується) однозначно вирішити питання про належність будь-якого елемента універсальної множини даній множині. Візьмемо, наприклад, довільний елемент a множини A і з'ясуємо, чи належить він множині C . У помилковому "визначенні" не сказано, із якої множини (з A чи B) складається множина C , тобто є невизначеність. Тому доти, поки не буде встановлено, із якої множини складається C , не можна буде вирішити питання про належність елемента a множині C . Якщо вибір буде зроблений на користь множини A , то $a \in C$, якщо ж буде обрана множина B , то $a \notin C$. У цьому й полягає помилковість даного "визначення". ▲

Дві множини A і B називаються **такими, що не перетинаються**, якщо вони не містять ніяких спільних елементів, тобто якщо їх перетин є порожня множина:

$$A \cap B = \emptyset.$$

Визначення множин, що не перетинаються, може бути поширене на більш ніж дві множини.

Множини $A_i, i = \overline{1, n}$ називаються такими, що **взаємно не перетинаються** (або **не перетинаються попарно**), якщо жодні дві з них не мають спільних елементів:

$$A_i A_j = \emptyset, \quad i, j = \overline{1, n}, \quad i \neq j.$$

Приклад 1.4 (продовження). A, B, C – взаємно неперетинні множини, тому що $AB = \emptyset, AC = \emptyset, BC = \emptyset$. ●

Доповненням множини A до множини B (зверніть увагу на порядок множин у записі дії !) називається множина E , що складається з усіх елементів B , які не належать A :

$$E = B \setminus A.$$

Доповненням множини A до універсальної множини U (або просто **доповненням множини A**) називається множина тих елементів U , які не належать A :

$$\overline{A} = U \setminus A.$$

Доповнення – це назва і дії над множинами, і її результату. Доповнення для множин A і B , показаних на рис. 1.1, пояснюється на рис. 1.4 (доповнення позначені сірим кольором).

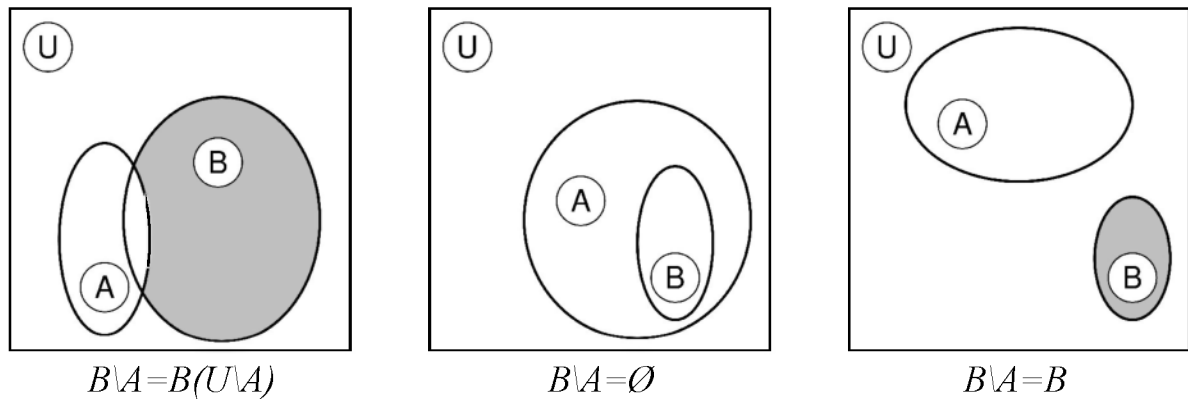


Рис. 1.4. Доповнення множин

Приклад 1.4 (продовження). $B \setminus A = B, \overline{A} = \{\text{люди, старші 10 років}\}$. ●

Зауваження 1.6. Для результату доповнення множин іноді використовується назва **різниця множин**, а для самої дії – таке ж позначення, що й для віднімання чисел:

$$E = B - A.$$

Легко переконатися, що "віднімання" множин не є дією, протилежною до додавання множин (див. рис. 1.2, 1.4):

$$A + (B - A) \neq A.$$

Тому далі буде використовуватися тільки введена раніше назва "доповнення". Утім, навіть у такої простої дії, як додавання множин, немає повної аналогії з додаванням чисел (розділ 1.4).

Ділення для множин не вводиться. ▲

Розбиттям множини A називається подання її у вигляді підмножин $A_i, i = \overline{1, n}$, які взаємно не перетинаються і в об'єднанні дають A :

$$A = \sum_{i=1}^n A_i = A_1 + A_2 + \dots + A_n, \quad A_i A_j = \emptyset, \quad i, j = \overline{1, n}, \quad i \neq j.$$

Розбиття – це назва дії, а результат не має спеціальної назви, оскільки розбиття утворюють більше ніж два об'єкти – множини. Нехай $U_1, U_2, \dots, U_n$ – розбиття універсальної множини U (рис. 1.5) і A – довільна підмножина U . Тоді розбиття U породжує й деяке розбиття підмножини A (рис. 1.6):

$$A = AU_1 + AU_2 + \dots + AU_n.$$

Для того щоб переконатися в цьому, необхідно перевірити, чи множини $AU_1, AU_2, \dots, AU_n$ взаємно не перетинаються й чи є їх об'єднання множиною A . Оскільки задане розбиття універсальної множини U , то деякий елемент $a \in A$ входить тільки в одну з підмножин $U_1, U_2, \dots, U_n$. Тоді елемент a входить тільки в одну із множин $AU_1, AU_2, \dots, AU_n$. Отже, ці множини попарно не перетинаються, а оскільки всі ці множини є підмножинами A , їх об'єднання дає A .

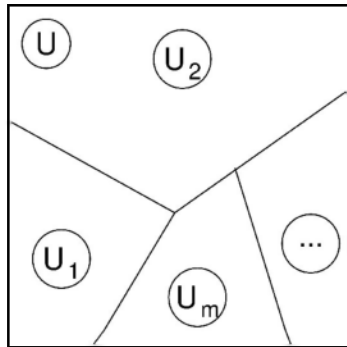


Рис. 1.5. Розбиття множини U довільними множинами

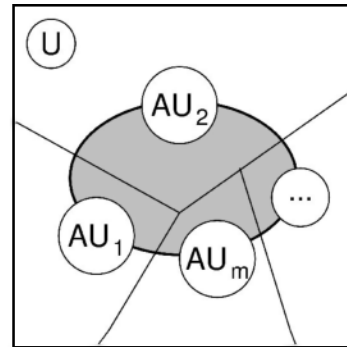


Рис. 1.6. Розбиття довільної множини A за рахунок розбиття множини U

Приклад 1.4 (продовження). Множини A, B, C задають розбиття U . Тоді будь-яка підмножина U буде мати очевидне розбиття, що задається множинами A, B, C . ●

Приклад 1.5. Нехай U – територія деякого континенту, тоді території держав $U_1, U_2, \dots, U_n$, розташованих на цьому континенті, задають розбиття U . Вони ж задають розбиття території A , зайнятої гірським масивом. ●

Зауваження 1.7. Часто зустрічається такий варіант "визначення" розбиття: "Розбиття множини – це розбиття її на суму частин, що не перетинаються". Будь-яке нове поняття **не можна** використовувати у визначенні його самого. Зверніть увагу на те, що в наведеному вище правильному визначенні не сказано нічого про те, як здійснене розбиття множини A . Ця інформація **не потрібна** для того, щоб з'ясувати, чи здійснює деяка група множин розбиття даної множини A . У визначенні сказано лише, що є множини, що не перетинаються, сума яких утворює множину A . ▲

Прямим добутком двох множин $A=\{a_1, a_2, \dots, a_n\}$ і $B=\{b_1, b_2, \dots, b_m\}$ називається множина $G=A \times B$ всіх можливих упорядкованих пар (двійок) елементів множин A і B :

$$G = \{(a_1, b_1), (a_1, b_2), \dots, (a_1, b_m), \\ (a_2, b_1), (a_2, b_2), \dots, (a_2, b_m), \\ \dots \\ (a_n, b_1), (a_n, b_2), \dots, (a_n, b_m)\}.$$

Прямий добуток – це назва і дії над множинами, і її результату.

Приклад 1.6. Нехай ген A має три алелі ($n=3$): $A=\{a_1, a_2, a_3\}$, а ген B – два ($m=2$): $B=\{b_1, b_2\}$. Тоді множина можливих генних наборів в одній неспареній хромосомі, що містить ці гени, може бути подана у вигляді прямого добутку (на рис. 1.7 упорядкованість елементів у парах показана стрілкою, спрямованою від елемента A до елемента B):

$$A \times B = \{a_1b_1, a_1b_2, a_2b_1, a_2b_2, a_3b_1, a_3b_2\}. \bullet$$

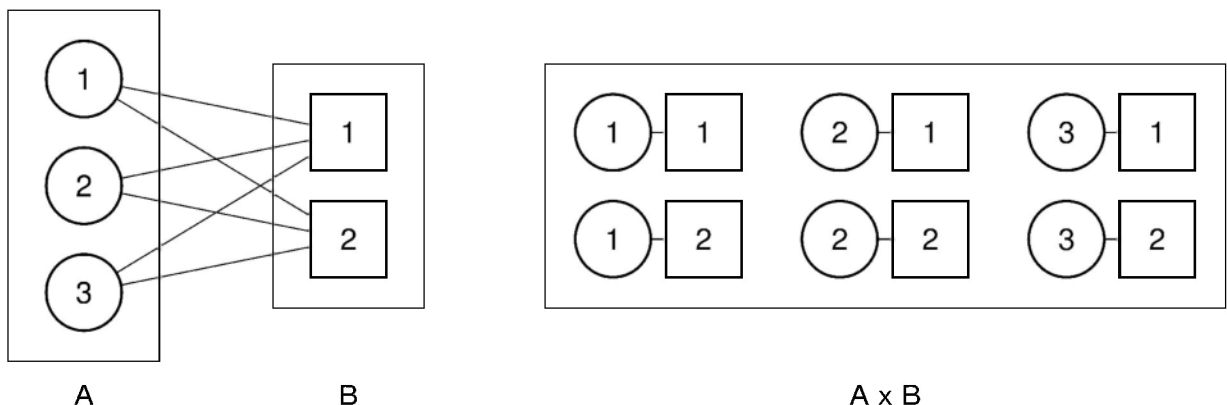


Рис. 1.7. Прямий добуток множин

Зауваження 1.8. Відмітна особливість прямого добутку множин порівняно з додаванням, множенням і доповненням множин полягає в тому, що структура множини-результату відрізняється від структур множин-операндів (так називаються множини, що беруть участь у деякій операції, тобто дії, наприклад, при додаванні множин A і B вони ж і є операндами). Щоб переконатися в цьому, достатньо подиви-

тися на рис. 1.7. Жоден з елементів прямого добутку двох множин $A \times B$ (жодна з упорядкованих пар $\circ \square$) не є елементом універсальних множин, із яких побудовані $A(\circ)$ і $B(\square)$. Тобто, в операндів прямого добутку можуть бути різні універсальні множини. ▲

Прямий добуток легко поширюється на довільну кількість множин. Зокрема, для трьох множин-співмножників A , B і C множина-результат містить упорядковані трійки елементів A , B і C .

Зауваження 1.9. Варто звернути увагу на різницю між відношеннями множин та дій над ними. Результатом відношень не є нова множина (так само, як і в разі порівняння двох чисел можна лише сказати, що одне з них більше іншого), тоді як у результаті дій утворюються нові множини (аналогічно тому, як при додаванні та інших операціях із числами утворюються нові числа – результати дій). ▲

1.4. Властивості дій над множинами

Додавання та множення множин мають властивості, аналогічні властивостям додавання й множення чисел (табл. 1.1).

Табл. 1.1

Властивість	Дія	
	Додавання	Множення
Переміщувальна	$A + B = B + A$	$A B = B A$
Сполучна	$(A + B) + C = A + (B + C)$	$(A B) C = A (B C)$
Розподільна	$A (B + C) = A B + A C$	

Зауваження 1.10. Властивості додавання й множення множин дозволяють оперувати виразами, у записі яких беруть участь множини-операнди та знаки “+” і “.”, так само, як і звичайними алгебраїчними виразами, тобто групувати члени, розкривати дужки й т. ін. Проте з цього **не випливає**, що є повна аналогія й відносно результатів дій – іноді вона існує і є прозора ($\emptyset$ і U виступають як аналоги 0 і ∞):

$$A + \emptyset = A; \quad A \emptyset = \emptyset, \quad A + U = U, \quad A U = U,$$

а іноді аналогія відсутня:

$$A + A = A; \quad A A = A.$$

Тому при роботі з довільними множинами не слід діяти за аналогією з діями над числами. Результати навіть найпростіших дій над множинами з погляду дій над числами можуть виглядати парадоксально. Наприклад, при додаванні у множину-результат не включаються **двічі** спільні елементи множин-операндів (якщо вони є), тобто чисельності множин не додаються. ▲

2. Елементи комбінаторики

2.1. Основні поняття комбінаторики

У теорії множин основна увага приділяється визначенню множин, діям над ними та співвідношенням між ними, причому кількість елементів тієї або іншої множини не має значення (див. зауваження 1.1). У комбінаториці, навпаки, робота із множинами базується на тому, що кількість елементів множин (підмножин деякої універсальної множини) має істотне значення. У комбінаториці виділяють три основних поняття – перестановки, сполучення й розміщення, які відповідають певним модельним ситуаціям – утворенням підмножин деякої універсальної множини. Ці поняття розглядаються в розділах 2.3, 2.4 і 2.5. Попередньо в розділах 2.2 і 2.3 вводяться принцип¹ математичної індукції й правило множення, які застосовуються для обґрунтування числових співвідношень, пов'язаних із перестановками, сполученнями й розміщеннями.

2.2. Принцип математичної індукції

Принцип математичної індукції² формулюється в такий спосіб.

Нехай $A(m)$ – вираз (формула), що залежить від натурального числа m . Якщо вираз $A(m_0)$ для $m_0=1$ правдивий (початок індукції) і з правдивості виразу для деякого m , тобто $A(m)$, випливає правдивість виразу для $m+1$, тобто $A(m+1)$ (крок індукції), то вираз $A(m)$ правдивий для всіх натуральних m .

Індукція може починатися не з $m_0=1$, а з будь-якого натурального числа $m_0 \geq 1$. У цьому випадку вираз $A(m)$ правдивий для всіх натуральних чисел m , більших m_0 .

2.3. Правило множення для множин

Нехай множина A_1 містить n_1 елементів, множина A_2 містить n_2 елементів, ..., множина A_m містить n_m елементів. Тоді кількість способів вибору по одному елементу з кожної множини дорівнює добутку

$$n_1 n_2 \dots n_m. \quad (2.1)$$

Приклад 1.6. (продовження). В одній неспареній хромосомі, що містить гени A і B , може бути $n \cdot m = 6$ генних наборів. ●

Приклад 2.1. (Гільдерман, 1974). Нехай множини A , B , C містять відповідно 2, 3 і 5 елементів. Тоді множина $A \times B \times C = \{(a, b, c) : a \in A, b \in B, c \in C\}$ містить $2 \times 3 \times 5 = 30$ елементів – трійок виду (a, b, c) . ●

¹ Від лат. *principium* (початок, основа) – основне базове положення деякої теорії, учення.

² Від лат. *inductio* (наведення) – умовивід від фактів до деякої гіпотези або загального твердження.

Доведемо правило множення для множин у загальному вигляді. Для цього зобразимо множини $A_1, A_2, \dots, A_m$ у вигляді коробок із кульками різних видів, причому кожна коробка містить кульки тільки одного виду, які для зручності пронумеровані. Перша коробка містить n_1 кульок, друга – n_2 , остання, m -на за числом, – n_m кульок, як показано на рис. 2.1.

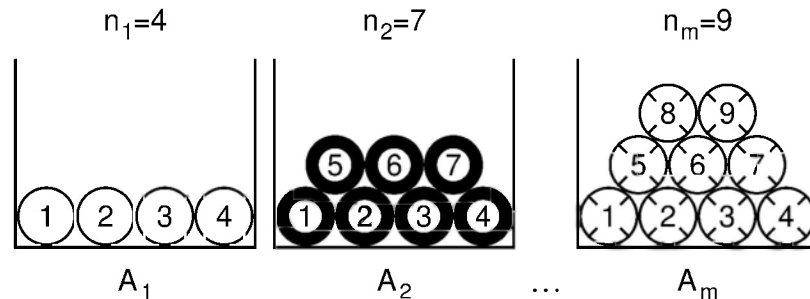


Рис. 2.1. Графічне пояснення правила множення (до початку індукції)

Перевіримо правило множення для однієї коробки (початок індукції). Очевидно, що з коробки, яка містить n_1 кульок, можна вибрати по одній кульці кількістю способів, що дорівнює n_1 . Тобто при $m=1$ правило множення “працює”.

Нехай $m = 2$ (цей крок уведений для наочності, але міг бути пропущений), тобто маємо дві коробки. Виберемо з першої коробки кульку з номером 1. У пару їй із другої коробки можна вибрати кульку з номером 1, потім – із номером 2, аж до останньої, з номером n_2 . Вийдуть пари, показані на рис. 2.2.

Усього пар, у яких одна кулька з номером 1 вибрана з першої коробки, а друга кулька з номером 1, 2, ..., n_2 – із другої коробки, вийде n_2 штук, за кількістю кульок у другій коробці.

Повернемо всі кульки у “свої” коробки й виберемо з першої коробки кульку з номером 2. До пари їй можна підібрати n_2 кульок із другої коробки (рис. 2.2).

Знову, як і в попередньому випадку, таких пар буде n_2 штук.

Продовжуючи складати пари кульок, одна з яких буде вибрана з першої коробки, а інша – із другої, дійдемо до останньої кульки з першої коробки, що має номер n_1 . Їй також можна підібрати до пари n_2 кульок із другої коробки. Кількість таких пар, як і раніше, дорівнює кількості n_2 кульок у другій коробці.

Підрахуємо кількість складених пар. Оскільки для кожної кульки з першої коробки кількість пар є однакою і дорівнює кількості кульок у другій коробці n_2 , а кількість кульок у першій коробці дорівнює n_1 , то загальна кількість пар дорівнює

$$\underbrace{n_2 + n_2 + \dots + n_2}_{n_1 \text{ разів}} = n_1 n_2.$$

Тобто правило множення для множин правдиве й для $m=2$.

Перейдемо до кроку індукції для довільного значення m . Це означає, що ми вважаємо правило множення правдивим для m коробок. Додамо до коробок ще одну, $(m+1)$ -ну за числом, у якій знаходиться n_{m+1} кульок “свого” виду (рис. 2.3).

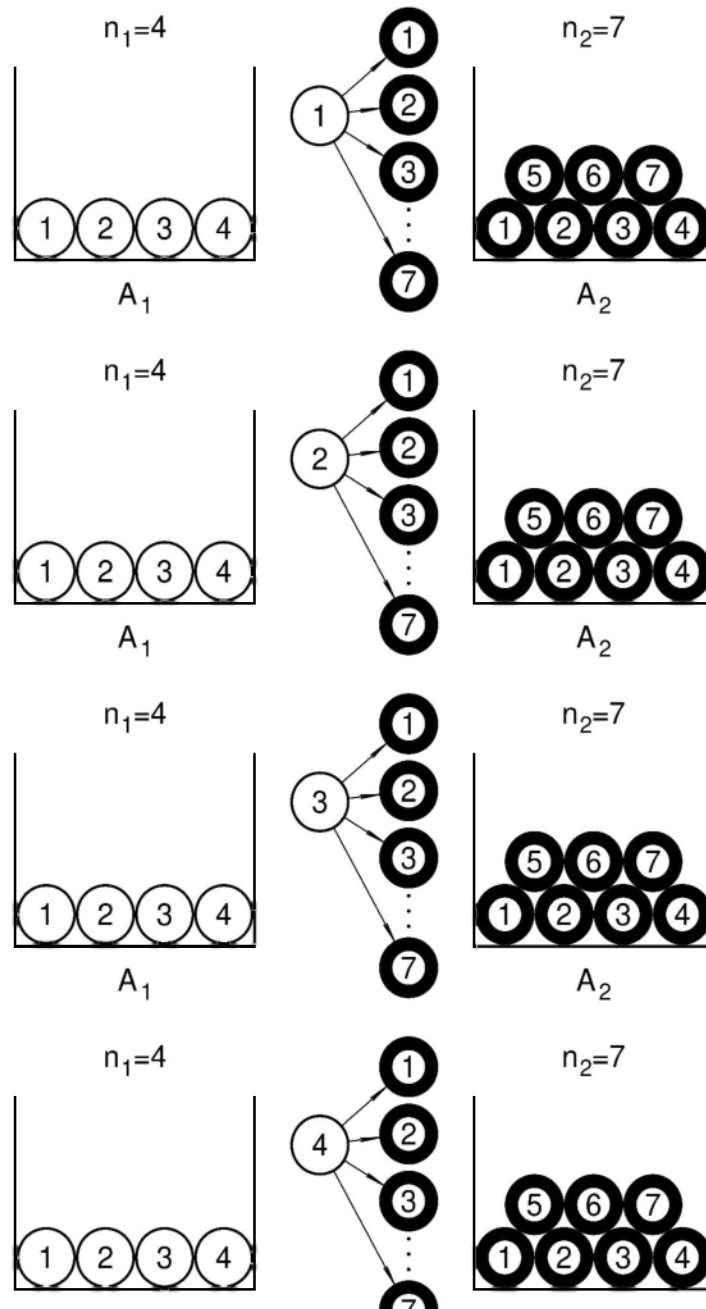


Рис. 2.2. Графічне пояснення правила множення (після початку індукції, $m=2$)

Покажемо, що із правдивості правила множення для m коробок впливає його правдивість для $(m+1)$ -ї коробки.

Виберемо з перших m коробок по одній кульці, у результаті вийде множина різних кульок – перша “ємка”. До цієї “ємки” додамо ще одну кульку з $(m+1)$ -ї коробки. Вийде множина з $(m+1)$ кульок. Оскільки в $(m+1)$ -ній коробці кількість кульок дорівнює n_{m+1} , то таких множин буде n_{m+1} .

Виконаємо такі ж дії з наступною, другою, “емкою”, вибраною по одній кульці з перших m коробок. До цієї “емки” знову можна додати одну кульку з $(m+1)$ -ї коробки. І знову вийде n_{m+1} множин, що складаються з $m+1$ кульок.

Так можна дійти до останньої “емки”, вибраної по одній кульці з перших m коробок. Оскільки на даному кроці індукції правило множення для m коробок вважається справедливим, то остання “емка” має номер, який дорівнює $n_1 n_2 \dots n_m$. Додаючи до цієї “емки” кульку з $(m+1)$ -ї коробки, знову одержимо n_{m+1} множин, що містять $m+1$ кульок.

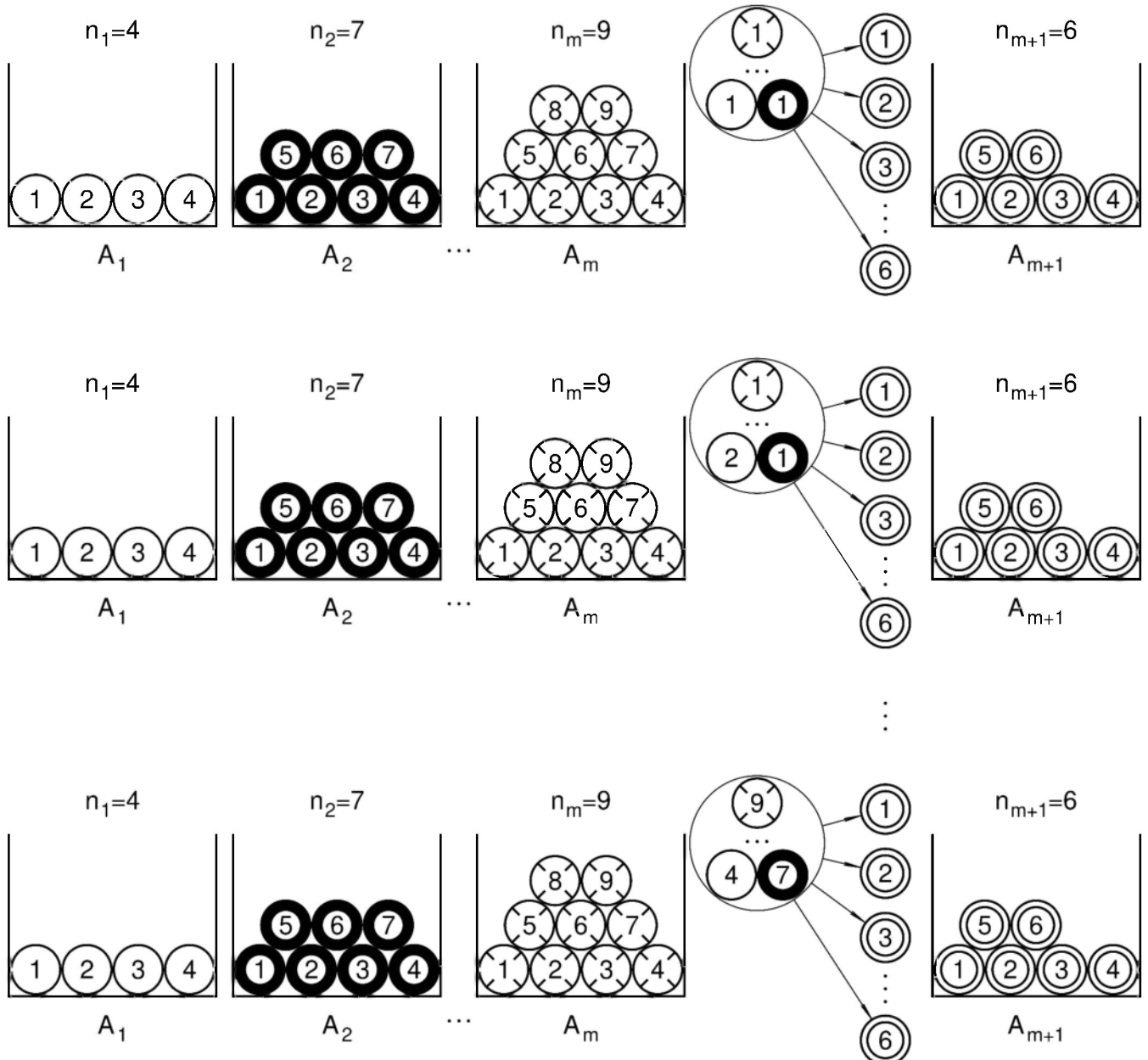


Рис. 2.3. Графічне пояснення правила множення (після початку індукції, $m \rightarrow m+1$)

Підведемо підсумки. Із кожної “емки” кульок, вибраної по одній кульці з перших m коробок, можна одержати множини, що складаються з $m+1$ кульок, додаючи до “емок” по одній кульці з додаткової $(m+1)$ -ї коробки. Оскільки кількість “емок” дорівнює $n_1 n_2 \dots n_m$, то кількість множин, які складаються з $m+1$ кульок, дорівнює

$$\underbrace{n_{m+1} + n_{m+1} + \dots + n_{m+1}}_{n_1 n_2 \dots n_m \quad \partial \dot{a} \zeta^3 \dot{a}} = (n_1 n_2 \dots n_m) n_{m+1}.$$

Таким чином, із правдивості правила множення для m множин випливає правдивість правила множення для $m+1$ множини. ■

2.4. Перестановки

Розглянемо деяку множину, що складається з m елементів, які розрізняються. Нехай такою множиною буде коробка з пронумерованими кульками (рис. 2.4).

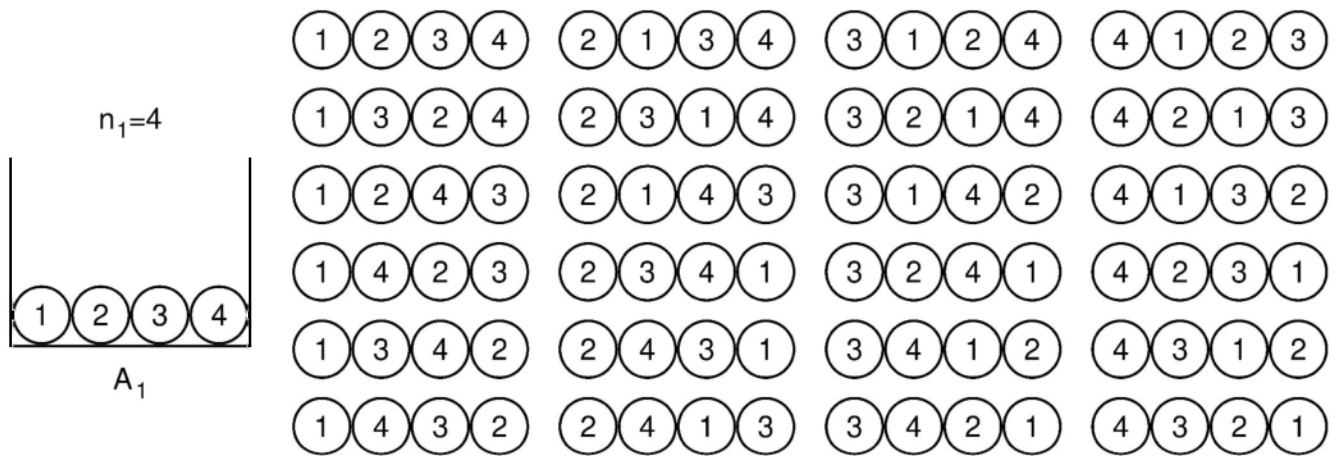


Рис. 2.4. Графічне пояснення визначення перестановок ($P_4 = 4! = 2 \cdot 3 \cdot 4$)

Будемо витягати по черзі всі кульки з даної коробки, викладаючи їх у певному порядку. Кожний такий порядок серед усіх елементів деякої множини будемо називати перестановкою. Очевидно, що для даної множини перестановка може бути організована не одним способом, і кількість перестановок певним чином пов'язана з кількістю елементів m множини. Перейдемо до формулювання строгих визначень і результатів, пов'язаних із перестановками.

Перестановка з m елементів є впорядкування цих елементів, тобто розташування m елементів у певному порядку.

Число перестановок з m елементів дорівнює

$$P_m = 1 \cdot 2 \cdot \dots \cdot m = m! \quad (2.2)$$

Справді, щоб розташувати m елементів у певному порядку, виберемо один елемент, що буде “першим”: зробити цей вибір можна m способами. Із $m-1$ елементів, що залишилися, виберемо другий елемент. Для цього є $m-1$ спосіб. Повторюючи цей процес k разів ($k=1, 2, \dots, m$), k -й елемент виберемо $m-k+1$ способами. Тоді відповідно до правила множення існує $m \cdot (m-1) \cdot (m-2) \cdot \dots \cdot 2 \cdot 1 = m!$ перестановок з m елементів. ■

Зауваження 2.1. Знак оклику після кількості елементів m у формулі (2.2) для кількості перестановок позначає цілу функцію цілого аргументу, що називається факторіалом. Читається запис $m!$ так: “факторіал ем”. ▲

Приклад 2.2. (Гільдерман, 1974). Для лікування хвороби застосовуються три засоби, і ефект лікування залежить від послідовності застосування засобів. Тоді можна призначити $P_3 = 6$ варіантів лікування: $\{1, 2, 3\}$, $\{1, 3, 2\}$, $\{2, 1, 3\}$, $\{2, 3, 1\}$, $\{3, 1, 2\}$, $\{3, 2, 1\}$. ●

Зауваження 2.2. Іноді доводиться розглядати “екзотичні” випадки множин, що складаються з $m=0$ елементів. Це порожні множини (див. розд. 1.1), елементи яких (відсутні!) потрібно впорядковувати. Такі множини прийнято впорядковувати єдиним способом, тобто $P_m = 0! = 1$. Це означає, що в такий спосіб функція факторіалу, визначена спочатку для натуральних чисел, визначається й для нуля, тобто $0! = 1$, а не 0.

2.5. Розміщення

Розглянемо деяку множину, що складається з m елементів, які розрізняються, наприклад коробку з кульками (див. розд. 2.4). Будемо витягати тепер не всі m елементів, а тільки їх частину із кількістю кульок, що дорівнює k , тобто підмножину з k елементів, але, як і раніше, серед витягнутих k елементів установлюється певний порядок (рис. 2.5).

Кожний такий порядок у підмножині вибраних елементів називається розміщенням. Перейдемо до строгих формулювань.

Розміщення з m елементів по k – будь-яка підмножина, що містить k елементів ($0 \leq k \leq m$), узятих у певному порядку з множини, що містить m елементів.

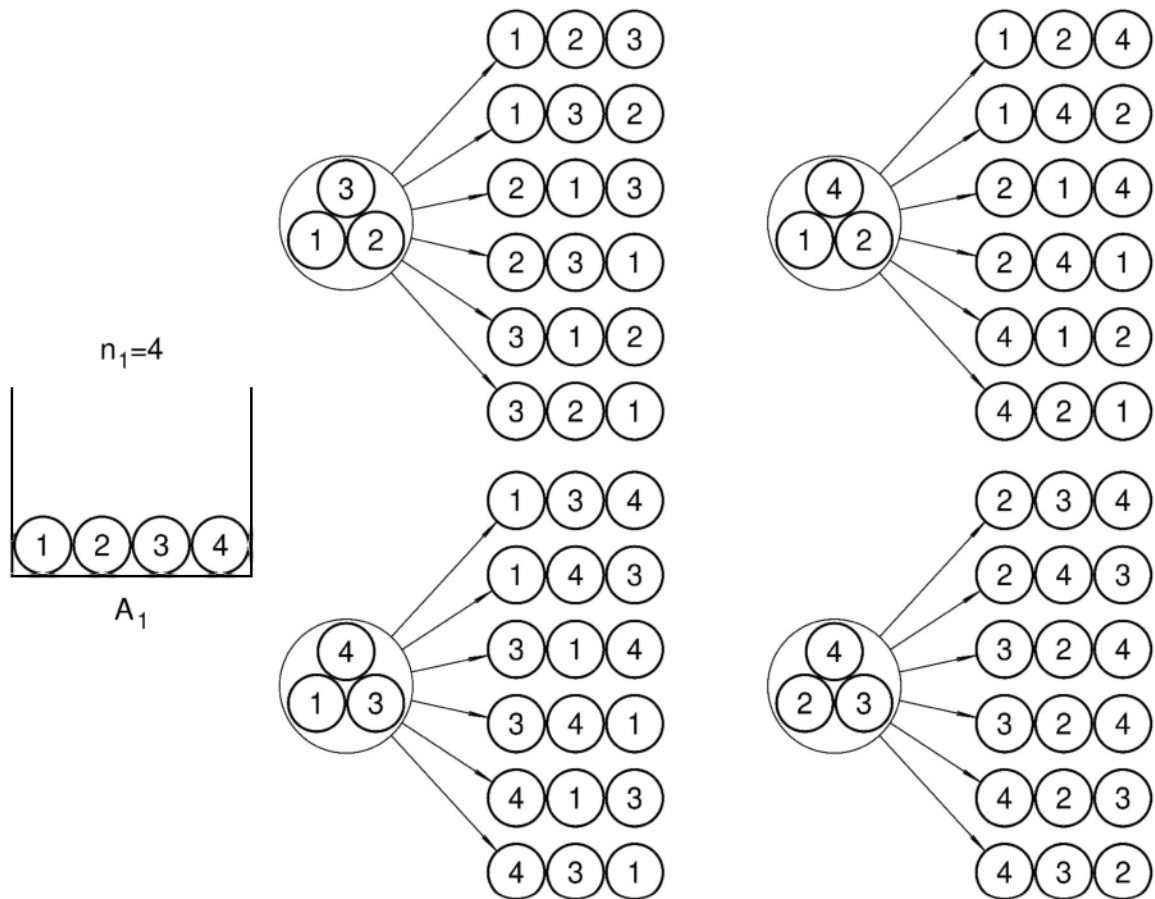
Кількість розміщень із m елементів по k дорівнює

$$A_m^k = \frac{m!}{(m-k)!}. \quad (2.3)$$

Дійсно, існує m способів вибору першого елемента, $m-1$ спосіб вибору другого елемента, $m-k+1$ способів вибору k -го елемента. За правилом множення маємо таке:

$$A_m^k = m \cdot (m-1) \cdot (m-2) \cdot \dots \cdot (m-k+1) = \frac{m \cdot (m-1) \cdot (m-2) \cdot \dots \cdot (m-k+1) \cdot (m-k) \cdot \dots \cdot 2 \cdot 1}{(m-k) \cdot \dots \cdot 2 \cdot 1} = \frac{m!}{(m-k)!}$$

Приклад 2.3. (Гільдерман, 1974). Нехай хвороба піддається лікуванню будь-якими двома засобами з відомих трьох і ефект лікування залежить від послідовності застосування засобів. Тоді можна призначити $A_3^2 = 6$ варіантів лікування: $\{1, 2\}$, $\{2, 1\}$, $\{1, 3\}$, $\{3, 1\}$, $\{2, 3\}$, $\{3, 2\}$. ●



міщень, установлювати не будемо. Витягнуту з m кульок підмножину k кульок без установленного порядку (тобто купку, жменьку й т.д.) будемо називати сполученням (рис. 2.6)

Розглянемо строгі визначення і результати, що стосуються сполучень.

Сполучення з m елементів по k – будь-яка підмножина, що містить k елементів, узятих із множини, що містить m елементів.

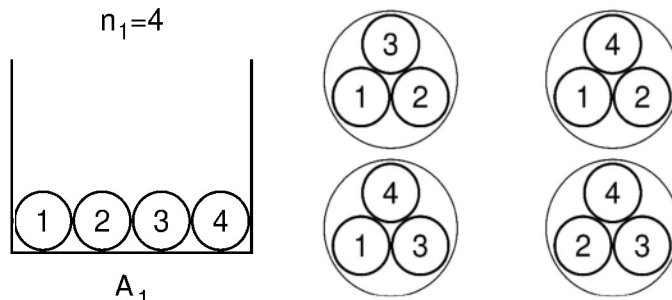


Рис. 2.6. Графічне пояснення визначення сполучень ($C_4^3 = 4$)

Кількість сполучень із m елементів по k дорівнює

$$C_m^k = \frac{m!}{k!(m-k)!} \quad (2.4)$$

Очевидно, що для одержання всіх можливих розміщень з m по k (їх кількість дорівнює A_m^k), достатньо взяти всі можливі сполучення з m по k (їх кількість дорівнює C_m^k), а потім із кожного сполучення утворити $k!$ можливих перестановок. Таким чином, $A_m^k = C_m^k k!$, звідки й випливає формула (2.4) для кількості сполучень. ■

Приклад 2.4. (Гільдерман, 1974). Нехай хвороба піддається лікуванню будь-якими двома засобами з відомих трьох, і ефект лікування не залежить від послідовності застосування засобів. Тоді можна призначити $C_3^2 = 3$ варіанти лікування: $\{1, 2\}$, $\{1, 3\}$, $\{2, 3\}$. ●

Зауваження 2.4. Формула (2.4) для кількості сполучень правдива й у граничних випадках $k=0$ і $k=m$ (див. зауваження 2.3). ▲

Зауваження 2.5. Поширена помилка полягає в тому, що перестановки, сполучення й розміщення ототожнюються з кількістю перестановок P_n , сполучень C_n^m і розміщень A_n^m . Слід пам'ятати, що перестановки, сполучення й розміщення – способи відбору й упорядкування елементів множини, а P_n , C_n^m і A_n^m – кількість таких способів. ▲

Зауваження 2.6. Правило, що дозволяє розрізняти перестановки, сполучення й розміщення, полягає в такому. Перестановка – це порядок, що встановлюється серед усіх елементів множини. Сполучення – довільна підмножина деякої множини, причому порядок, у якому елементи вибирають у підмножину, не має значення. Розміщення відрізняється від сполучення тим, що порядок вибору має значення. ▲

3. Випадкові події

3.1. Причинність і випадковість у природознавстві

Відношення між **причинністю** і **випадковістю** завжди були одним із важливих питань у вивченні природи. У будь-якому явищі загальна кількість складників, або елементів, що взаємодіють між собою, є нескінченно велика, тому всі зв'язки між ними врахувати неможливо. Однак головною характеристикою більшості явищ є обмежена кількість елементів⁴, і їх взаємодія може бути реально простежена. Тоді стає зрозуміло, які взаємодії яких елементів виступають причиною явища, а які виявляються наслідком причини. При цьому причина завжди передує наслідку, і винятків тут бути не може. Наслідок завжди з'являється після причини й ніколи – до неї. Якщо таке нібито трапляється, то це означає, що існує інша причина, а можливо, неявний ланцюг причин, що й викликає цей наслідок, а те, що зовні здається причиною, насправді нею не є.

Наявність деякої невизначеності в ланцюзі причинно-наслідкових відношень свідчить про неповноту дослідження причин, адже причина ніколи не існує в чистому вигляді, вона супроводжується іншими прихованими причинами або умовами. Оскільки в будь-якому явищі неможливо врахувати всі елементи, що взаємодіють між собою, то в наслідку можуть бути виявлені відхилення від тих результатів, до яких приводять основні причинні фактори. Про такі відхилення прийнято говорити, що вони випадкові. Насправді вони є наслідком неврахованих факторів, що впливають на загальний результат, і завдання дослідника полягає у відокремленні одних причин від інших і встановленні закономірних причинно-наслідкових відношень у явищах.

Приклад 3.1. (Вентцель, 1970). Монета підкидається певним чином, після чого фіксується результат – сторона, звернена догори ("орел" або "решка"). За найбільш ретельного дотримання умов підкидання при наступному підкиданні результат не визначений. ●

Приклад 3.2. (Вентцель, 1970). Один і той же предмет кілька разів зважується на дуже точних (аналітичних) вагах, однак результати повторних зважувань незначно відрізняються один від одного. ●

Цілком очевидно, що в природі немає жодного явища, у якому б не були присутні тією чи іншою мірою елементи випадковості. Як би точно не фіксувалися умови досліду, неможливо досягти того, щоб у разі його повторення результати повністю й точно збігалися. Але так вважалося не завжди.

У картині світу, побудованій класичною наукою в період з XVII по XVIII ст., будь-яке явище визначалося початковими умовами, що задаються (принаймні в принципі) абсолютно точно². У такому світі не було місця випадковості, всі деталі його були ретельно приладжені й знаходились "у зчепленні", подібно до шес-

⁴ Виділення таких елементів загалом є досить складне завдання. Різні дослідники в тих самих умовах можуть вирішити цю проблему по-різному.

² Зміна модельних поглядів на Всесвіт викладена згідно з ідеями І. Пригожина (1987).

терень якоїсь космічної машини. Поширення механістичного світогляду збіглося в часі з розквітом машинної цивілізації. У машинне століття Всесвіт зображувався як гігантський механізм. Саме механістичний погляд лежить в основі знаменитого вислову Лапласа про те, що істота, здатна охопити всю сукупність даних про стан Всесвіту в будь-який момент часу, могла б не тільки точно пророчити майбутнє, але й до дрібних деталей відновити минуле (**демон Лапласа**). Уявлення про простий і однорідний Всесвіт не тільки зробило вирішальний вплив на хід розвитку науки, але й залишило помітний відбиток на інших сферах людської діяльності.

Багато слабких сторін механістичної моделі було виявлено досить давно. Уявлення про світ як про годинниковий механізм із планетами, що споконвіку обертаються по незмінних орбітах, із детермінованою поведінкою будь-яких рівноважних систем і діючими на всі без винятку об'єкти універсальними законами, які можуть бути відкриті зовнішнім спостерігачем, – така модель із самого початку зазнавала нищівної критики. Це пов'язане з тим, що в такому світі неможлива еволюція, ідею якої висунув Дарвін, тобто ідея розвитку стає зовсім недоречною.

На початку ХІХ ст. термодинаміка поставила під сумнів позачасовий характер механістичної картини світу. Якби світ був гігантською машиною, то така машина неминуче повинна була б зупинитися, оскільки запас корисної енергії рано або пізно був би вичерпаний. Світові годинники не могли б іти вічно.

Незабаром після цього послідовники Дарвіна висунули протилежну ідею. На їх думку, хоча світова машина, витрачаючи енергію й переходячи з більш організованого в менш організований стан, і могла сповільнювати свій хід і навіть зупинятися, проте біологічні системи повинні були розвиватися тільки по висхідній лінії, переходячи з менш організованого в більш організований стан.

Трохи пізніше фізики, які працювали в галузі квантової механіки, знову поставили під сумнів правдивість детерміністичної моделі, принаймні на рівні мікросвіту. Однак і на рівні макросвіту прикладів, що вказували на набагато важливішу роль випадковості, знаходилося усе більше.

Приклад 3.3. (Пригожин, 1986). Мала кулька під дією деякого поштовху може зіштовхуватися з декількома довільним чином розташованими нерухомими кулями (рис. 3.1). Траєкторія малої кульки цілком визначена (детермінована). Але досить у початкові умови руху малої кульки внести невелику невизначеність, як у результаті послідовних зіткнень початкова невизначеність підсилиться й унеможливить простеження положення кульки. ●

У цьому прикладі цілком детермінована механістична система виявилася нестійкою (слабко детермінованою): досить малі розходження в початкових умовах із часом підсилюються.

Сучасний науковий погляд такий. Деякі частини Всесвіту справді можуть діяти як механізми. Це замкнуті системи, але вони становлять лише малу частку фізичного Всесвіту. Більшість же систем є відкриті, вони обмінюються енергією або речовиною з навколишнім середовищем. До відкритих систем належать біологічні й соціальні системи, а це означає, що будь-які спроби зрозуміти їх у рамках механістичної моделі приречені на невдачу.

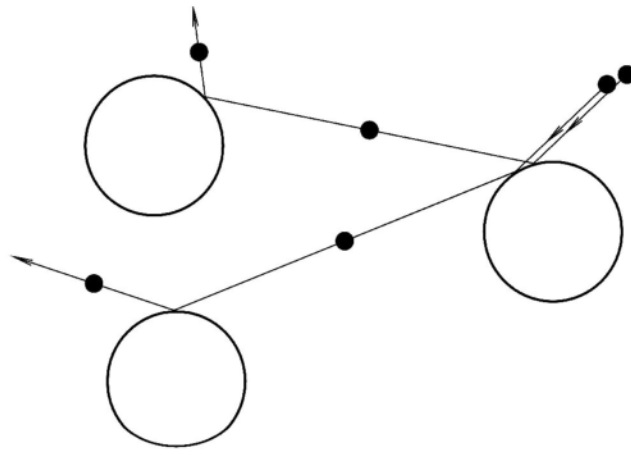


Рис. 3.1. Графічне пояснення слабкої детермінованості

Крім того, відкритий характер переважної більшості систем у Всесвіті наводить на думку про те, що **головну роль у навколишньому світі відіграють нестійкість і нерівноважність**. Перед нестійкими системами, можливо, і демон Лапласа виявився б не всемогутнім.

3.2. Поняття досліду і події

Дослідом називається відтворювана сукупність умов, за яких спостерігається певний результат.

Досліди, у яких результат незмінно повторюється, називаються дослідом з **детермінованим** (тобто визначеним або не випадковим) **результатом**.

Досліди, у яких результат не повторюється, називають дослідом з **недетермінованим або випадковим** (тобто невизначеним) **результатом**.

Результат досліду має спеціальну назву – **подія**.

Подія, що спостерігається в досліді з випадковим результатом, називається **випадковою**.

Зауваження 3.1. Надалі для нас не буде мати великого значення причина, через яку результат досліду не змінюється. Виявлення причин випадковості або не випадковості того чи іншого явища – справа конкретної предметної галузі природознавства, у якій працює дослідник. Ми можемо виходити з погляду, який назовемо моделлю “чорного ящика”. У цій моделі “чорний ящик” має деякий запас відгуків на зовнішній вплив, а яким буде відгук, дослідник заздалегідь не знає, тобто детермінований причинно-наслідковий ланцюжок відсутній. За таких умов потрібно виявити закономірності функціонування “чорного ящика”. Зауважимо також, що при зверненні до випадкової події ми маємо на увазі якісний результат. Такий результат фіксується в лабораторному журналі не за допомогою кількісних показників, а записами відносно того, що щось відбулося (наприклад, “+”) або не відбулося (“-”). ▲

Події позначаються великими літерами латинської абетки. Якщо слід указати, у чому саме полягає настання тієї чи іншої події, після літери, що вказує на

подію, ставлять знак рівності “=”, за яким у фігурних дужках у той чи інший спосіб розкривають зміст події.

3.3. Ймовірність випадкової події

Ймовірність випадкової події A – це її числова характеристика, що показує ступінь можливості появи цієї події. Ймовірність випадкової події A позначається через $P(A)$.

Зауваження 3.2. Важлива не сама по собі ймовірність окремо взятої випадкової події, а можливість порівняння ймовірностей двох або декількох подій, одна з яких досить добре відома й може бути еталоном. Так, якщо ймовірність події A дорівнює 0,3, то це ще ні про що не говорить. Але якщо в тих же умовах ймовірність іншої події B дорівнює 0,4, то можна вважати, що умови для появи другої події більш сприятливі або друга подія спостерігається частіше першої (див. розділ 3.4). Проте мало вміти порівнювати ймовірності декількох подій, потрібно вміти їх визначати (в одних дослідах визначення ймовірності нагадує, скоріше, обчислення за формулою, в інших – вимір фізичної величини). ●

Перш ніж вводити той або інший спосіб визначення ймовірності, розглянемо як приклад фізичну аналогію, що ґрунтується на введенні поняття температури й способу виміру температури. Інтуїтивне визначення температури як ступеня нагріву, що відчувається людиною, досить суб'єктивне й навряд чи може бути сприйняте іншим. Тому при введенні способу виміру температури домовляються про два таких поняття:

1) реперні, або опорні, точки на шкалі (їх дві), тобто такі стани, які можуть бути легко відтворені за їх описами і яким зовсім довільно приписуються певні значення температури;

2) фізичний процес, за допомогою якого буде проводитися вимір температури; як правило, мова йде про зміну об'єму так званого “термометричного тіла”, наприклад, ртуті, спирту і т. ін.

Приклад 3.4. У шкалі Цельсія як реперні точки обрані стани замерзання води (танення льоду) і кипіння води (конденсації пари) при стандартному атмосферному тиску. Цим реперним точкам приписані значення температури, які дорівнюють відповідно 0°C і 100°C . На довільність цих числових значень указує той факт, що спочатку Цельсій приписав таненню льоду температуру 100°C , а кипінню води – 0°C . Відомі нам значення температури узвичаїлися трохи пізніше, уже після Цельсія. Для того щоб звичайний процес зміни об'єму “термометричного тіла” при його нагріванні або охолодженні міг бути застосований для виміру температури, потрібно ввести шкалу, що, по суті, являє собою лінійку, на якій між двома точками із значеннями 0 і 100 є 99 поділок. Певне положення межі термометричного тіла, наприклад меніска підфарбованого спирту або ртуті, відповідає певній поділці на шкалі, а отже, значенню температури. ▲

Повернемося до визначення ймовірності, почавши з домовленості про вибір двох реперних (опорних) подій. Такими реперними подіями вибирають неможливу та вірогідну події.

Неможливою подією називають подію, що в умовах даного досліді ніколи не відбувається.

Вірогідною називається подія, яка в умовах даного досліді відбувається завжди.

Подія, що не є неможлива або вірогідна, називається (і є) **можливою**.

За визначенням, неможливій події приписують ймовірність, яка дорівнює 0, а вірогідній – ймовірність, яка дорівнює 1. Ймовірність можливих подій знаходиться між 0 і 1:

$$0 < P(A) < 1.$$

Тобто різниця між ймовірностями вірогідної й неможливої подій відіграє роль масштабу для виміру ймовірностей усіх інших подій. Ніщо не заважає приписати опорним подіям інші значення ймовірності, наприклад -100 і $+100$, але склалася традиція вибору саме значень 0 і 1.

Якщо подія, що цікавить дослідника, наприклад A , у даному досліді не настала, то настала якась інша подія. Таких інших подій у даному досліді може бути досить багато, причому всі вони для дослідника, що вивчає саме подію A^3 , можуть бути нецікавими. Для того щоб відзначити в лабораторному журналі, що подія A не настала, не вдаючись до зайвих подробиць стосовно того, що ж усе-таки відбулося в даному досліді, вводять спеціальне визначення.

Протилежною до події A називається подія $\bar{A}$, що полягає в ненастанні події A .

Якщо якась подія A є неможливою, то протилежна їй подія $\bar{A}$ буде вірогідною і навпаки.

Деяка подія в умовах даного досліді називається **практично неможливою**, якщо її ймовірність можна вважати дуже малою:

$$P(A) \approx 0.$$

Деяка подія в умовах даного досліді називається **практично вірогідною**, якщо її ймовірність можна вважати близькою до одиниці:

$$P(A) \approx 1.$$

Якщо подія A практично неможлива, то протилежна їй подія $\bar{A}$ буде практично вірогідною і навпаки.

Зауваження 3.3. У багатьох людей, які починають вивчати теорію ймовірностей, виникає питання, яке саме значення ймовірності варто вважати малим, а яке – великим. Відповідь на це питання досить суб'єктивна, і не може бути дана точно. У теорії ймовірностей і математичній статистиці досить часто зустрічають-

³ Можливо, що заради вивчення закономірності настання тільки події A даний досліді і проводиться.

ся ситуації, які викликають подив у “новачка”, тому що в них явно присутній суб’єктивізм дослідника. Частково роз’яснює ситуацію поданий нижче приклад. ▲

Приклад 3.5. (Вентцель, 1970). У початківців виникає питання, чи можна вважати малим значення ймовірності, яке дорівнює 10^{-3} . Припустимо, що отримана позитивна відповідь. Тоді виникає інше запитання, що є уточненням першого: чи можна вважати малою ймовірність 10^{-3} того, що під час стрибка з парашутом парашут не розкриється? Звичайно, відразу ж дається негативна відповідь. ●

На завершення даного розділу сформулюємо один принцип, який разом із поняттями практично неможливої й практично вірогідної події стане остаточно зрозумілим і потрібним при вивченні основ статистики. Нагадаємо, що принцип – це таке твердження, що не доводиться й, загалом кажучи, не може бути доведено. Однак це твердження досить обґрунтоване практично, не призводить до протиріч і застосовується разом з іншими подібними твердженнями як методологічна база певної науки (можна сказати, що такі твердження, або домовленості, складають **парадигму науки**).

Принцип практичної вірогідності. Якщо в умовах даного дослідження подія A практично неможлива, то при одноразовому проведенні дослідження подію A можна вважати взагалі неможливою.

Очевидність цього принципу не вимагає пояснень. Наприклад, купуючи лотерейний квиток, людина, яка раціонально мислить, не буде брати великий кредит і т. ін., сподіваючись на виграш.

3.4. Статистична ймовірність

Нехай дослід проводиться n разів (будемо вважати, що довжина серії дослідів дорівнює n або що проведена серія дослідів довжиною n). Нехай у даній серії дослідів наставала подія A . Кількість настань (появ, спостережень і т. ін.) події A в серії n дослідів називається **частотою події A** . Частота події позначається $n(A)$. Сама частота події $n(A)$ безвідносно до довжини n серії дослідів мало про що говорить. Дійсно, нехай подія A настає 10 разів у серії зі 100 дослідів, а подія B – стільки ж разів у серії з 1000 дослідів. Частоти подій A і B однакові, тобто $n(A) = n(B)$, але різниця в настанні подій A і B очевидна. Для того щоб “прив’язати” частоту до довжини серії дослідів, простіше розглянути відношення цих двох величин, тобто ввести нову величину.

Відносна частота події A – це відношення частоти події A до довжини серії дослідів:

$$\frac{n(A)}{n}.$$

Оскільки частота події не може бути меншою нуля й більшою n , то значення відносної частоти будь-якої події знаходиться в межах

$$0 \leq \frac{n(A)}{n} \leq 1.$$

Життєвий досвід і здоровий глузд підказують, що при збільшенні кількості дослідів для деяких подій відносна частота перестає з прийнятною для дослідника точністю змінюватися. У таких випадках говорять, що в умовах даного дослідів для відносної частоти виконується **властивість стабілізації**.

Приклад 3.6. (Вентцель, 1970). Нехай дослід полягає в підкиданні монети, а подія, що цікавить нас, визначена так: $A = \{\text{поява орла}\}$. У табл. 3.1 наведені результати серії дослідів з $n=600$ підкидань монети (для простоти дослід був поділений на 60 “підсерій”, у кожній із яких кидалися одночасно 10 монет).

Табл. 3.1

N	$P^*(A)$	N	$P^*(A)$	N	$P^*(A)$
10	0,600	210	0,462	410	0,502
20	0,650	220	0,472	420	0,512
30	0,600	230	0,470	430	0,512
40	0,575	240	0,479	440	0,514
50	0,540	250	0,484	450	0,519
60	0,550	260	0,477	460	0,515
70	0,528	270	0,489	470	0,515
80	0,512	280	0,482	480	0,510
90	0,588	290	0,493	490	0,508
100	0,490	300	0,497	500	0,510
110	0,550	310	0,500	510	0,506
120	0,492	320	0,503	520	0,516
130	0,523	330	0,497	530	0,513
140	0,500	340	0,506	540	0,515
150	0,493	350	0,497	550	0,513
160	0,475	360	0,497	560	0,511
170	0,471	370	0,495	570	0,512
180	0,472	380	0,492	580	0,509
190	0,463	390	0,500	590	0,507
200	0,465	400	0,498	600	0,505

Для більшої наочності на рис. 3.2 зображена залежність відносної частоти появи орла від кількості дослідів n . ●

Зауваження 3.4. Зрозуміло, що зовсім не обов'язково в кожній серії дослідів буде спостерігатися стабілізація відносної частоти. Можливо, для розглянутих процесів характерна наявність деякого “збурюючого” фактора, що непомітно для дослідника змінює умови дослідів. Досить уявити підкидання монети, яку після кожного підкидання гнуть. Такі процеси часто зустрічаються в природознавстві, але нами розглядатися не будуть, тому що умови дослідів тут змінюються. ▲

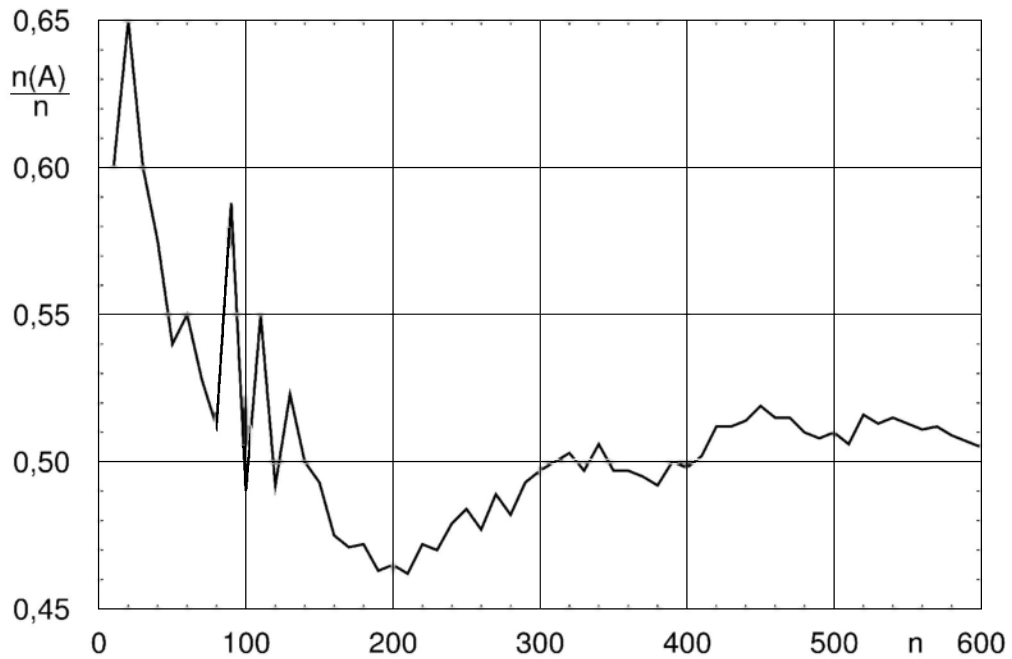


Рис. 3.2. Графік стабілізації відносних частот

Якщо в досліді спостерігається стабілізація відносної частоти, будемо називати стабілізовану відносну частоту **статистичною ймовірністю** й позначати верхнім індексом у вигляді “зірочки”⁴

$$P^*(A) = \frac{n(A)}{n}. \quad (3.1)$$

Іноді, щоб підкреслити важливість виконання умови стабілізації відносної частоти визначення статистичної ймовірності записують у такий спосіб:

$$P^*(A) = \lim_{n \rightarrow \infty} \frac{n(A)}{n}. \quad (3.2)$$

Очевидно, що статистична ймовірність задовольняє загальне обмеження для значень ймовірності будь-якої події:

$$0 \leq P^*(A) \leq 1. \quad (3.3)$$

Зауваження 3.5. Такий експериментальний вимір ймовірності може викликати деякий подив. Однак відзначимо, що багато фізичних вимірів саме так і здійснюється. Наприклад, при вимірі температури вимога стабілізації у вигляді вирівнювання температур фізичного тіла й термометра (його теплоємність повинна бути набагато меншою теплоємності тіла, температура якого вимірюється) є не тіль-

⁴ Зірочкою тут і далі будемо позначати величини, визначені в досліді.

ки звичайна, вона до того ж формулюється у вигляді так званого нульового початку термодинаміки. Без можливості такого вирівнювання (стабілізації) експериментальна основа термодинаміки просто не могла б бути забезпечена. ▲

3.5. Класична ймовірність

Для введення іншого способу визначення поняття ймовірності випадкових подій потрібні кілька допоміжних визначень, зокрема такі.

Кілька подій в умовах даного досліду називаються **повною групою (подій)**, якщо дослід завжди (можна сказати, обов'язково) закінчується однією з цих подій.

Приклад 3.7. Якщо дослід полягає в підкиданні грального кубика, то повну групу можна скласти, перелічивши всі результати даного досліду: $\{1\}$, $\{2\}$, $\{3\}$, $\{4\}$, $\{5\}$, $\{6\}$. ●

Повна група подій має таку властивість, що називається **властивістю повноти**: якщо до повної групи подій додати довільну кількість можливих у даному досліді подій, то властивість повноти групи подій не зміниться.

Приклад 3.7 (продовження). Так, якщо додати до списку з шести подій ще дві, які полягають у випаданні парного й непарного числа: $B=\{2,4,6\}$ і $C=\{1,3,5\}$, то можна переконатися, що властивість повноти “розширеної” групи подій, як і раніше, виконується. ●

Кілька подій в умовах даного досліду називаються **несумісними**, якщо поява однієї з них виключає настання інших (у досліді може з'явитися тільки одна з несумісних подій).

Приклад 3.7 (продовження). Шість подій $\{1\}$, $\{2\}$, $\{3\}$, $\{4\}$, $\{5\}$, $\{6\}$ утворюють групу несумісних подій. Для того щоб не склалося помилкове уявлення, що повна група й група несумісних подій – це одне й те саме, відразу вкажемо, що додавання події B або C робить цю групу подій сумісною. Дійсно, щоразу, коли настає подія $\{1\}$, настає й подія C . ●

У групи несумісних подій, так само, як і в повної групи подій, є одна очевидна властивість. Вона не має ніякої спеціальної назви й формулюється так. Із групи несумісних подій можна виключати довільну кількість подій доти, поки кількість подій, що залишилися, буде не менша двох, при цьому властивість несумісності в групі подій, що залишилися, зберігається.

Кілька подій в умовах даного досліду називаються **рівноможливими**, якщо дослід організований таким чином, що заздалегідь відомо, що жодна з подій не має більшої можливості появи, ніж інші.

Приклад 3.7. (продовження). Якщо гральний кубик симетричний, то події $\{1\}$, $\{2\}$, $\{3\}$, $\{4\}$, $\{5\}$, $\{6\}$ утворюють групу рівноможливих подій. ●

На основі визначень понять повної групи, групи несумісних подій і групи рівноможливих подій дамо нове визначення. Воно буде застосоване для визначення класичної ймовірності. Отже, повна група несумісних і рівноможливих подій в умовах даного досліду називається **випадками (шансами)**. Усі випадки, що можуть спостерігатися в даному досліді, розділимо на дві категорії (на два сорти, два різновиди й т.д.) стосовно їх відношень до події A , що цікавить нас.

Випадок називається **сприятливим** для події A , якщо настання даного випадку викликає (можна сказати, означає, причому автоматично) настання події A .

Приклад 3.7 (продовження). Для події B сприятливими будуть такі події: $\{2\}$, $\{4\}$, $\{6\}$, а для події C – три події, що залишилися з групи несумісних подій: $\{1\}$, $\{3\}$, $\{5\}$. ●

Нехай в умовах даного дослідження можуть спостерігатися m випадків, із яких лише $m(A)$ сприяють події A , тоді **класичною ймовірністю** події A називається відношення⁵:

$$P(A) = \frac{m(A)}{m}. \quad (3.4)$$

Оскільки кількість сприятливих даній події випадків $m(A)$ не може бути меншою нуля й більшою кількості всіх випадків m даного дослідження, то значення класичної ймовірності будь-якої події знаходиться в межах

$$0 \leq P(A) \leq 1, \quad (3.5)$$

що узгоджується із загальним визначенням ймовірності, даним у розд. 3.3.

Приклад 3.8 (Вентцель, 1970). У коробці є 5 куль, із яких 2 білі і 3 чорні. Із коробки навмання виймають одну кулю. Знайти ймовірність того, що ця куля буде білою.

Для події $A = \{\text{поява білої кулі}\}$ загальна кількість випадків $m=5$; із них $m(A) = 2$, тобто 2 випадки є сприятливі події A , отже $P(A) = 2/5 = 0,4$. ●

Приклад 3.9 (Вентцель, 1970). У коробці є 7 куль: 4 білі і 3 чорні. З неї виймають (одночасно або послідовно) дві кулі. Знайти ймовірність того, що обидві кулі будуть білими.

Для події $B = \{\text{обидві кулі білі}\}$ загальна кількість випадків дорівнює кількості способів, якими можна вибрати 2 кулі з 7:

$$m = C_7^2 = \frac{7!}{(7-2)!} = \frac{6 \cdot 7}{1 \cdot 2} = 21,$$

а кількість сприятливих події B випадків – це кількість способів, якими можна вибрати 2 білі кулі з 4:

$$n(B) = C_4^2 = \frac{3 \cdot 4}{1 \cdot 2} = 6.$$

⁵ Ми навмисно відступаємо від традиції, відповідно до якої прийнято говорити про “класичну формулу ймовірності”, причому тільки лише для однаковості з визначенням ймовірності, даним у розділі 3.4. Це, на нашу думку, сприятиме більш легкому запам'ятовуванню. Втім, можна було б зовсім “симетрично” говорити про “статистичну формулу ймовірності”.

Звідси $P(B)=6/21=2/7$. ●

Приклад 3.10 (Вентцель, 1970). У коробці 3 білі й 2 чорні кулі. Визначимо події:

$A = \{ \text{поява білої кулі при першому вийманні} \},$

$B = \{ \text{поява білої кулі при другому вийманні} \}.$

Якщо перша вийнята куля повертається в коробку, тоді, очевидно, імовірність настання події B не залежить від того, чи настала подія A :

$$P(A)=3/5=0,6, \quad P(B)=3/5=0,6.$$

Якщо ж перша вийнята куля не повертається в коробку, то ситуація змінюється. Якщо подія A настала, тобто була вийнята біла куля, то ймовірність події B дорівнює $P(B) = 1/2$. Якщо ж подія A не настала, тобто була вийнята чорна куля, то $P(B) = 3/4$. Таким чином, ймовірність події B набуває різних значень залежно від того, настала чи не настала подія A . У таких випадках говорять, що **подія B залежить від події A** (тобто настання події B залежить від настання події A), а ймовірність події називають **умовною** (див. розділ 3.13). ●

3.6. Порівняння статистичної та класичної ймовірностей

При всій подібності формул для статистичної (3.1) і класичної (3.4) ймовірностей (зрозуміло, що букви m і n можна поміняти місцями у формулах, від цього нічого не зміниться; головне – домовитися про те, яка буква що позначає) ці поняття принципово розрізняються. Звернемо увагу на те, що статистична ймовірність може обчислюватися для будь-якої події. Потрібні тільки:

1) можливість багаторазового відтворення відповідного досліду (ця умова “закладена” у визначенні досліду);

2) виконання умови стабілізації відносної частоти.

Сфера застосування класичної ймовірності набагато вужча, оскільки умови досліду повинні бути такими, щоб його результати задовольняли визначення випадків (шансів). Для цього потрібна деяка “симетрія” (не обов’язково в буквальному, геометричному, сенсі) умов досліду. Такі умови не дуже часто зустрічаються в природознавстві. Наприклад, їх можна спостерігати в генетиці й квантовій механіці. Важливо інше. У тих дослідах, для яких застосування класичної ймовірності випадкових подій є можливе, виявляється, що класична й статистична ймовірності приводять до одних і тих же значень. Це свідчить про те, що введені нами способи визначення ймовірності узгоджені між собою, що дозволяє сподіватися на їх об’єктивність.

3.7. Геометрична ймовірність

У багатьох дослідах кількість подій така велика, що можна вважати її практично нескінченною. У подібних випадках використання класичного визначення ймовірності є неприйнятне, і для її обчислення застосовують геометричну інтерпретацію (Гільдерман, 1991).

Усі можливі результати таких дослідів зображуються точками деякого відрізка U на прямій, точками деякої фігури U на площині (рис. 3.3) або точками деякого тіла в просторі, а сприятливі деякій події A результати – точками відповідної фігури A – підмножини множини U . За визначенням ймовірність події A , якій сприяють зазначені події, дорівнює

$$P(A) = \frac{|A|}{|U|}, \quad (3.6)$$

де знак $| |$ позначає не модуль, а довжину, площу чи об'єм (загальна назва – міра) відповідної фігури.

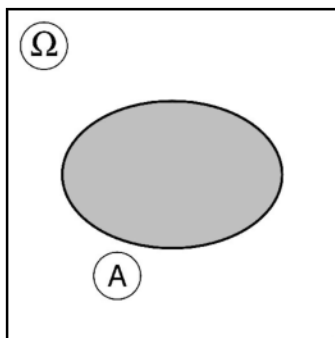


Рис. 3.3. Зображення випадкової події на діаграмі Венна

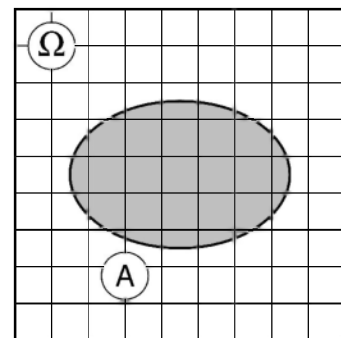


Рис. 3.4. Ілюстрація геометричної ймовірності випадкової події

Геометричне визначення ймовірності пов'язане з класичним. Справа в тому, що фігуру U на площині можна розбити на дрібні квадратики (рис. 3.4), і кожній випадковій події можна буде поставити у відповідність не точку, а такий квадратик. Тоді кількість усіх рівноможливих подій – це кількість усіх квадратиків, що складають фігуру U , тобто величина, пропорційна площі U , а кількість подій, які сприяють події A , – це кількість квадратиків, у фігурі A , тобто величина, пропорційна площі фігури A . Звідси й виходить формула (3.6) для геометричної ймовірності.

Зрозуміло, що кількість квадратиків дає лише наближене значення площі фігури A і, для того, щоб одержати точне значення, потрібно перейти до границі відношення кількості квадратиків, що складають фігуру A , до кількості квадратиків, які утворюють фігуру U , за необмеженого зменшення розмірів квадратиків (тобто за необмеженого збільшення їх кількості). Зміст геометричного визначення й полягає в тому, що не треба щоразу розбивати фігури на квадратики, обчислювати їх

розмір, визначати їх кількість і переходити до границі. Достатньо знайти площі фігур A і U і обчислити відношення їх площ.

3.8. Простір елементарних подій

У розд. 3.1 – 3.6 ми ввели важливі поняття, пов'язані з імовірністю випадкових подій. Звернемо увагу на те, що ми розглядали тільки поодинокі події. Тим часом у природознавстві велике значення має вміння виявляти взаємозв'язок між різними, найчастіше випадковими, явищами. Розглянемо випадкові події, дії над ними та зв'язки між ними. Попередньо наведемо визначення деяких понять.

Деяка подія в умовах даного досліду називається **елементарною**, якщо вона не може бути зведена до інших, більш простих в умовах даного досліду, подій.

Приклад 3.11. У досліді підкидається гральний кубик. Подія $B = \{\text{випадіння парного числа}\}$ не є елементарна, тому що може бути зведена до таких подій: $\{2\}$, $\{4\}$, $\{6\}$. Останні, у свою чергу, не можуть бути зведені до яких-небудь інших подій, а тому є елементарні. ●

Множина всіх елементарних подій в умовах даного досліду називається **простором елементарних подій**. Простір елементарних подій позначається великою літерою Ω грецької абетки, а для позначення елементарної події використовується мала літера ω , причому за необхідності з індексом.

Установимо відповідність основних понять теорії множин і теорії ймовірностей (табл. 3.2), що буде далі використовуватися для наочності доведень.

Табл. 3.2.

Теорія множин	Теорія ймовірностей
Універсальна множина U	Простір елементарних подій Ω , вірогідна подія
Множина A	Подія A
Елемент множини a	Елементарна подія ω
Порожня множина $\emptyset$	Неможлива подія

Із погляду теорії множин довільна подія в умовах даного досліду є відповідною підмножиною простору елементарних подій.

Для зручності розгляду відношень між подіями та дій над ними можна використовувати діаграми Венна. У разі використання діаграм Венна для зображення подій, відношень між ними та дій над ними зручною виявляється геометрична інтерпретація ймовірності (див. розд. 3.7). Ймовірність $P(A)$ події A обчислюється як відношення площ фігур, що зображують події A і Ω (рис. 3.3):

$$P(A) = \frac{S(A)}{S(\Omega)} \quad (3.7)$$

3.9. Дії над подіями

Демо визначення дій над подіями двома способами. Перший спосіб ґрунтується на відомій термінології теорії множин, а другий – більше відповідає теорії ймовірностей. Можна сказати інакше: перший спосіб підкреслює “склад” події, тобто вказує на елементарні події, що її утворюють, а другий – на спостережуваний у досліді результат. Наведемо обидва визначення: перше, “мовою” теорії множин, – ліворуч, а друге, ймовірнісне, – праворуч.

Сумою двох подій A і B називається подія

$$A + B = C,$$

яка складається з елементарних подій, що належать або до події A , або до події B (хоч би до однієї з цих подій).	яка полягає в настанні або події A , або події B (хоча б однієї з цих подій).
---	---

Сума є результатом дії над подіями, що називається додаванням і позначається знаком “+”.

Приклад 3.11 (продовження). Для подій $A = \{1, 2, 3\}$ і $B = \{2, 4, 6\}$ сумою буде подія $A + B = C = \{1, 2, 3, 4, 6\}$. ●

Добутком двох подій A і B називається подія

$$A \cdot B = A B = D,$$

яка складається з елементарних подій, що належать і до події A , і до події B (до обох подій).	яка полягає в настанні і події A , і події B (в одночасній появі цих подій).
--	--

Добуток є результатом дії над подіями, що називається множенням і позначається знаком “.” або взагалі не позначається.

Приклад 3.11 (продовження). $A B = D = \{2\}$. ●

Доповненням події A до події B називається подія

$$E = B \setminus A,$$

яка складається з елементарних подій, що належать до події B , але не належать до події A .	яка полягає в настанні події B , але не в настанні події A .
---	--

Доповнення є одночасно назвою і результату дії над подіями, і самої дії, воно позначається так : “\”.

Приклад 3.11 (продовження). $B \setminus A = E = \{4, 6\}$. ●

Інколи трапляються ситуації, коли деяка подія A доповнюється до простору елементарних подій, тобто коли $B = \Omega$. Для результату такого доповнення вводиться спеціальна назва.

Подією, протилежною даній події A , називається подія

$$\bar{A} = \Omega \setminus A,$$

яка складається з елементів Ω , що не належать до події A . | яка полягає в ненастанні події A .

Звичайно про події A і $\bar{A}$ говорять як про **прямі** та **протилежні** події.

Уведемо відношення двох подій на основі добутку. Дві події називають **несумісними**, якщо їх спільне настання є подією неможливою, тобто $A \cap B = \emptyset$.

Якщо подія A може бути подана у вигляді суми m несумісних подій $A_1, A_2, \dots, A_m$

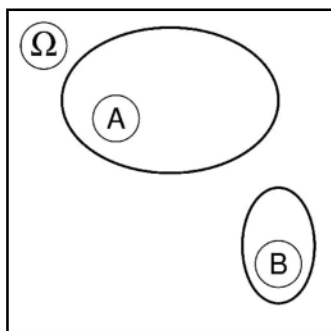
$$A = A_1 + A_2 + \dots + A_m, \quad A_i A_j = \emptyset, \quad i, j = \overline{1, m}, \quad i \neq j,$$

то події $A_1, A_2, \dots, A_m$ називаються **варіантами** події A .

3.10. Правило додавання ймовірностей для двох несумісних подій

Розглянемо дві несумісні події A і B , ймовірності яких $P(A)$ і $P(B)$ відомі. Поставимо задачу визначення ймовірності $P(A+B)$ суми $A+B$ цих подій.

Розв'яжемо цю задачу двома способами. Перший спосіб ґрунтується на використанні діаграм Венна та геометричній інтерпретації ймовірності, а другий – на статистичній ймовірності. Викладаючи другий спосіб, будемо вважати, що відомі результати серії з n дослідів, подані у вигляді “лабораторного журналу”, у якому наведені дані про настання (“плюс”) та ненастання (“мінус”) подій A і B (умова стабілізації відносних частот вважається виконаною). Обидва розв'язання подамо одночасно: перше, “геометричне”, – ліворуч, а друге, статистичне, – праворуч.



№	A	B
1	+	–
2	–	+
3	–	–
...	...	...
n	...	...

“Лабораторний журнал” при спостереженні двох несумісних подій може

Оскільки A і B – несумісні події, то події $A+B$ відповідає фігура, що складається з двох частин A і B , які не перетинаються. Отже площа $S(A+B)$ фігури $A+B$ дорівнює сумі площ $S(A)$ і $S(B)$ фігур A і B . Тому

$$P(A+B) = \frac{S(A+B)}{S(\Omega)} = \frac{S(A) + S(B)}{S(\Omega)} = \frac{S(A)}{S(\Omega)} + \frac{S(B)}{S(\Omega)},$$

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

містити тільки рядки з двома “мінусами” або одним “мінусом” і одним “плюсом”.

Частота події $A+B$ дорівнює сумі частот подій A і B . Можна сказати інакше: для підрахунку частоти події $A+B$ можна зробити підрахунок частот окремо за стовпчиками A і B , а потім скласти ці частоти. Оскільки в одному рядку двох “плюсів” бути не може, то частота $n(A+B)$ буде обчислена правильно. Тому в разі виконання умови стабілізації відносної частоти для статистичної ймовірності маємо

$$P^*(A+B) = \frac{n(A+B)}{n} = \frac{n(A)}{n} + \frac{n(B)}{n},$$

$$P^*(A+B) = P^*(A) + P^*(B).$$

Сформулюємо одержаний результат у вигляді такого правила: **ймовірність суми двох несумісних подій дорівнює сумі ймовірностей цих подій:**

$$P(A+B) = P(A) + P(B). \quad (3.8)$$

Якщо як подію B вибрати подію $\bar{A}$, протилежну даній події A , то правило додавання для несумісних подій можна буде записати таким чином:

$$P(A + \bar{A}) = P(A) + P(\bar{A}). \quad (3.9)$$

Оскільки з визначення події $\bar{A}$ випливає, що $A + \bar{A} = \Omega$, то

$$P(A + \bar{A}) = P(\Omega) = 1. \quad (3.10)$$

Звідси маємо **перший висновок** із правила додавання ймовірностей для двох несумісних подій:

$$P(A) + P(\bar{A}) = 1, \quad (3.11)$$

який можна сформулювати так: **сума ймовірностей прямої та протилежної подій дорівнює одиниці.**

На практиці інтерес викликає ситуація, коли протилежна подія розкладається на меншу, порівняно з прямою подією, кількість варіантів. У такій ситуації ймові-

рність протилежної події обчислюється, як правило, простіше ймовірності прямої події. На цьому базується **другий висновок** із правила додавання ймовірностей для двох несумісних подій, що називається також **правилом переходу до протилежної події** і формулюється як алгоритм обчислення ймовірності прямої події:

- 1) визначається ймовірність протилежної події $P(\bar{A})$;
- 2) обчислюється ймовірність прямої події $P(A) = 1 - P(\bar{A})$.

На те, що протилежна подія $\bar{A}$ розкладається на меншу, порівняно з прямою подією, кількість варіантів, указують уточнення “принаймні”, “хоча б” у формулюванні задачі.

3.11. Правило додавання ймовірностей для декількох несумісних подій

Поширимо правило додавання ймовірностей для несумісних подій (3.8) на m несумісних подій $A_1, A_2, \dots, A_m$.

Подамо суму m подій у вигляді суми двох подій A і B :

$$A = A_1 + A_2 + \dots + A_{m-1}, \quad B = A_m,$$

тобто

$$A = A_1 + A_2 + \dots + A_{m-1} + A_m = A + B.$$

Тоді для двох несумісних подій A і B можна застосувати правило (3.8):

$$P(A_1 + A_2 + \dots + A_{m-1} + A_m) = P(A) + P(B) = P(A_1 + A_2 + \dots + A_{m-1}) + P(A_m).$$

Виконуючи послідовно такі ж самі перетворення з іншими подіями $A_1, A_2, \dots, A_m$, одержимо таке правило:

$$P(A_1 + A_2 + \dots + A_m) = P(A_1) + P(A_2) + \dots + P(A_m) = \sum_{i=1}^m P(A_i), \quad (3.12)$$

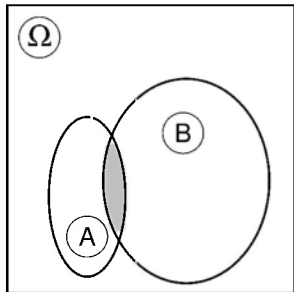
яке формулюється в такий спосіб: **ймовірність суми декількох несумісних подій дорівнює сумі ймовірностей цих подій.**

3.12. Правило додавання ймовірностей для двох довільних подій

Розв'яжемо поставлену в розд. 3.10 задачу на обчислення ймовірності суми двох подій для двох довільних подій A і B (вони можуть бути і сумісними, і несумісними). Розглянемо можливі ситуації відношення подій A і B , а розв'язання проведемо двома способами (ліворуч – геометричне розв'язання, праворуч – ста-

тистичне). Умова стабілізації відносних частот для даних у “лабораторному журналі” вважається виконаною.

1) $A \cap B \neq \emptyset$



$$\begin{aligned}
 P(A + B) &= \frac{S(A + B)}{S(\Omega)} = \\
 &= \frac{S(A) + S(B) - S(AB)}{S(\Omega)} = \\
 &= \frac{S(A)}{S(\Omega)} + \frac{S(B)}{S(\Omega)} - \frac{S(AB)}{S(\Omega)} = \\
 &= P(A) + P(B) - P(AB)
 \end{aligned}$$

“Лабораторний журнал” може містити рядки з чотирма комбінаціями “плюсів” і “мінусів”:

№	A	B
1	+	–
2	–	+
3	+	+
4	–	–
...	...	...
n	...	...

Для того щоб підрахувати частоту $n(A+B)$ події $A+B$, можна виконати підрахунки окремо за стовпчиками, тобто підрахувати частоти $n(A)$ і $n(B)$ подій A і B . Однак при цьому будуть двічі підраховані рядки з двома “плюсами”, тобто ті, у яких спостерігалася подія AB . Кількість таких рядків дорівнює частоті $n(AB)$ події AB . Тому для частоти події $A+B$ одержуємо вираз

$$n(A+B) = n(A) + n(B) - n(AB).$$

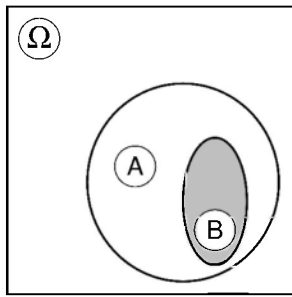
Для статистичної ймовірності маємо

$$P^*(A + B) = \frac{n(A + B)}{n} = \frac{n(A) + n(B) - n(AB)}{n},$$

$$P^*(A + B) = P^*(A) + P^*(B) - P^*(BA).$$

2) $B \subset A$

Щоразу, коли в “лабораторному журналі” у стовпчику B з’являється “плюс”, у стовпчику A також з’являється “плюс”. Із настання події A не впливає настання події B , тому “плюсу” в стовпчику A може відпові-



Оскільки події $A+B$ відповідає фігура A , то

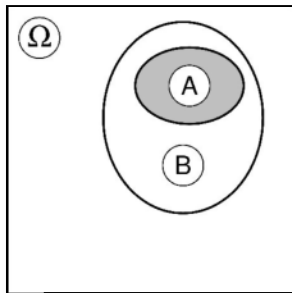
$$S(A+B)=S(A),$$

Звідси

$$P(A+B) = \frac{S(A+B)}{S(\Omega)} = \frac{S(A)}{S(\Omega)},$$

тобто

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



Оскільки події $A+B$ відповідає фігура B , то

$$S(A+B)=S(B),$$

звідки

$$P(A+B) = \frac{S(A+B)}{S(\Omega)} = \frac{S(B)}{S(\Omega)},$$

тобто $P(A+B) = P(B).$

дати “мінус” у стовпчику B . Обидві події можуть одночасно не відбуватися, тоді у відповідному рядку з'являються два “мінуси”.

№	A	B
1	+	—
2	+	+
3	—	—
...	...	...
n	...	...

Зрозуміло, що в процесі підрахунку кількості “плюсів” у стовпчику A ми не пропустимо жодного випадку настання події $A+B$, тому

$$n(A+B) = n(A).$$

Отже, для статистичної ймовірності одержуємо вираз

$$P^*(A+B) = \frac{n(A+B)}{n} = \frac{n(A)}{n} = P^*(A).$$

3) $A \subset B$

“Лабораторний журнал” виглядає так само, як у попередній ситуації, якщо поміняти події A і B місцями:

№	A	B
1	—	+
2	+	+
3	—	—
...	...	...
n	...	...

Отже, для частоти події $A+B$ маємо вираз

$$n(A+B) = n(B),$$

а для статистичної ймовірності

$$P^*(A+B) = \frac{n(A+B)}{n} = \frac{n(B)}{n} = P^*(B).$$

Найбільш загальний вигляд має формула, одержана для ситуації 1) сумісних подій. Перевіримо, чи дійсно вона охоплює всі розглянуті ситуації, включаючи розглянуту в розд. 3.10 ситуацію для несумісних подій:

а) оскільки $B \subset A$, то $P(AB)=P(B)$ і

$$P(A+B) = P(A) + P(B) - P(AB) = P(A) + P(B) - P(B) = P(A);$$

б) оскільки $A \subset B$, то $P(AB)=P(A)$ і

$$P(A+B) = P(A) + P(B) - P(AB) = P(A) + P(B) - P(A) = P(B);$$

в) для несумісних подій $P(AB)=0$, тому

$$P(A+B) = P(A) + P(B) - P(AB) = P(A) + P(B).$$

Отже, формула, одержана для ситуації сумісних подій, дійсно охоплює всі розглянуті ситуації й може вважатися найбільш загальною. На її основі сформулюємо правило додавання ймовірностей для двох довільних подій: **ймовірність суми довільних подій дорівнює сумі ймовірностей цих подій мінус ймовірність одночасного настання обох подій:**

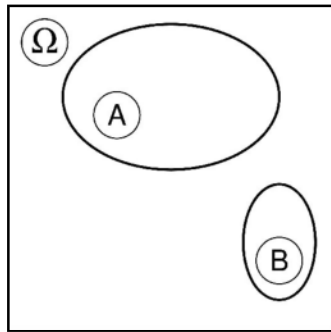
$$P(A+B) = P(A) + P(B) - P(AB). \quad (3.13)$$

3.13. Умовна ймовірність

Розглянуті раніше способи визначення ймовірності випадкової події ґрунтувалися тільки на “властивостях” події, що розглядається. Така ймовірність іноді називається безумовною. На практиці часто доводиться розглядати настання деякої події у зв'язку з настанням іншої події (див. приклад 3.10 у розд. 3.5). Через це вводиться ще одне визначення ймовірності, яку називають умовною. Спочатку розглянемо змістове визначення. Воно не пов'язане прямо з деякою формулою, а тільки вказує, що саме змінюється у визначенні умовної ймовірності порівняно з безумовною.

Умовною ймовірністю події A при настанні події B називається ймовірність події A , обчислена будь-яким способом з урахуванням настання події B . Позначається умовна ймовірність так: $P(A | B)$.

Розглянемо ситуації для подій A і B і спробуємо застосувати подане визначення для обчислення умовної ймовірності. Як і в, у розд. 3.10 і 3.12, обчислення виконаємо двома способами (умова стабілізації відносних частот для подій A і B у “лабораторному журналі” вважається виконаною).

1) $AB = \emptyset$ 

Оскільки після настання події B події A буде відповідати порожня множина $A|B = \emptyset$, то площа відповідної фігури дорівнюватиме нулю, отже

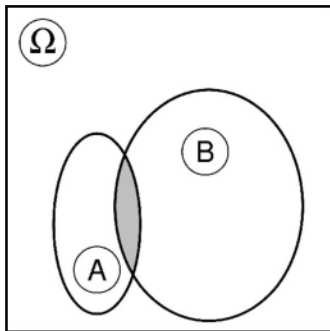
$$P(A|B) = 0.$$

У “лабораторному журналі” для двох несумісних подій в одному рядку не можуть з’явитися два “плюси”: (див. розд. 4.3):

№	A	B
1	+	–
2	–	+
3	–	–
...	...	...
n	...	...

Отже, підрахунок частоти події A за рядками, у стовпчику B яких знаходиться “плюс”, дасть нульове значення умовної частоти $n(A|B) = 0$. Тому

$$P^*(A|B) = 0.$$

2) $AB \neq \emptyset$ 

Після настання події B настання події A може відбутися тільки за рахунок тих елементарних подій B , які є спільними для A і B , тобто за рахунок елементарних подій, що належать AB . Тому для визначення умовної ймовірності події A при настанні B одержуємо

“Лабораторний журнал” може містити рядки з чотирма комбінаціями “плюсів” і “мінусів”:

№	A	B
1	+	–
2	–	+
3	+	+
4	–	–
...	...	...
n	...	...

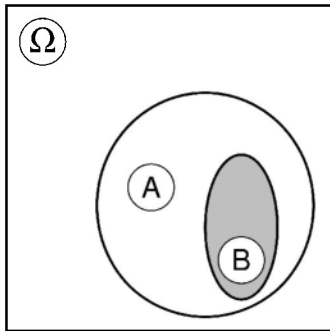
Підрахунок частоти події A після настання події B можна здійснити так: підраховуємо спочатку кількість рядків, у стовпчику B яких знаходиться “плюс”. Кількість таких рядків дорівнює частоті $n(B)$ події B . Потім серед цих рядків підраховуємо кількість рядків, у стовпчику A яких знаходиться “плюс”. Це число дорівнює частоті $n(AB)$ події AB . За ви-

$$P(A|B) = \frac{S(AB)}{S(B)} = \frac{S(AB)/S(\Omega)}{S(B)/S(\Omega)} = \frac{P(AB)}{P(B)}.$$

значенням статистичної ймовірності

$$P^*(A|B) = \frac{n(AB)}{n(B)} = \frac{n(AB)/n}{n(B)/n} = \frac{P^*(AB)}{P^*(B)}.$$

3) $B \subset A$



Після настання події B автоматично настає подія A , тому

$$P(A|B) = 1.$$

Щоразу, коли в “лабораторному журналі” у стовпчику B з’являється “плюс”, у стовпчику A також з’являється “плюс”.

№	A	B
1	+	–
2	+	+
3	–	–
...	...	...
n	...	...

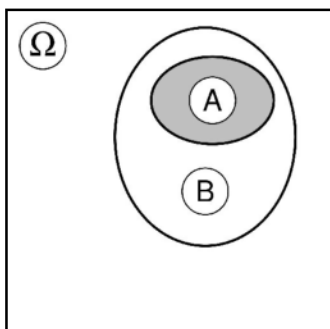
Зрозуміло, що підрахувавши кількість “плюсів” у стовпчику B , ми одержимо значення умовної частоти події A :

$$n(A|B) = n(B).$$

Тому вираз для умовної статистичної ймовірності матимемо такий вигляд:

$$P^*(A|B) = \frac{n(A|B)}{n(B)} = \frac{n(B)}{n(B)} = 1.$$

4) $A \subset B$



№	A	B
1	–	+
2	+	+
3	–	–
...	...	...
n	...	...

“Лабораторний журнал” виглядає

Після настання події B можливі результати досліду звужуються: зрозуміло, що умовну ймовірність події A потрібно визначати, порівнюючи площу фігури, що зображує подію A , і площу фігури, що зображує подію B , а не фігури, що зображує вірогідну подію Ω :

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

так само, як у попередній ситуації, якщо поміняти події A і B місцями.

Зрозуміло, що в результаті підрахунку кількості “плюсів” у стовпчику A ми одержимо значення умовної частоти події A :

$$n(A|B) = n(A).$$

Тому для умовної статистичної ймовірності отримаємо такий вираз:

$$P^*(A|B) = \frac{n(A|B)}{n(B)} = \frac{n(A)}{n(B)} = \frac{n(A)/n}{n(B)/n},$$

$$P^*(A|B) = \frac{P(A)}{P(B)}.$$

Найбільш загальний вигляд має формула, одержана для ситуації 2). Перевіримо, чи охоплює вона інші ситуації:

а) оскільки $A \cap B = \emptyset$, то $P(AB) = 0$ і

$$P(A|B) = \frac{P(AB)}{P(B)} = \frac{0}{P(B)} = 0;$$

б) оскільки $B \subset A = \emptyset$, то $AB = B$, $P(AB) = P(B)$ і

$$P(A|B) = \frac{P(AB)}{P(B)} = \frac{P(B)}{P(B)} = 1;$$

в) оскільки $A \subset B = \emptyset$, то $AB = A$, $P(AB) = P(A)$ і

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

Отже, формула, одержана для ситуації 2), дійсно, охоплює всі розглянуті ситуації й може вважатися найбільш загальною.

Дамо тепер інше визначення умовної ймовірності. **Умовною ймовірністю події A при настанні події B** називається величина, яка дорівнює відношенню ймовірності спільного настання подій A і B і ймовірності події B :

$$P(A|B) = \frac{P(AB)}{P(B)}. \quad (3.14)$$

Зазначимо, що у визначенні умовної ймовірності події A при настанні події B остання може бути тільки можливою, тобто $P(B) \neq 0$.

3.14. Правило множення ймовірностей для двох довільних подій

Визначення умовної ймовірності (3.14) застосовується також і самостійно у вигляді правила множення для двох подій:

$$P(AB) = P(B)P(A|B) = P(A)P(B|A), \quad (3.15)$$

яке формулюється в симетричному щодо подій A і B вигляді: **ймовірність спільного настання двох подій дорівнює добутку безумовної ймовірності однієї з них і умовної ймовірності другої в разі настання першої.**

3.15. Правило множення ймовірностей для декількох довільних подій

Поширимо правило множення ймовірностей для двох довільних подій (3.15) на m довільних подій

$$A_1, A_2, A_3, \dots, A_m.$$

Для обчислення ймовірності спільного настання цих подій подамо добуток m співмножників $A_1 A_2 A_3 \dots A_{m-1} A_m$ у вигляді добутку двох множників A і B :

$$A = A_1 A_2 A_3 \dots A_{m-1}, \quad B = A_m$$

тобто

$$(A_1 A_2 A_3 \dots A_{m-1})(A_m) = AB.$$

Тоді до двох подій A і B можна застосувати правило (3.15):

$$P(A_1 A_2 \dots A_{m-1} A_m) = P(A_m)P(A_m | A_1 A_2 \dots A_{m-1}).$$

Виконуючи послідовно такі перетворення з іншими подіями $A_1, A_2, \dots, A_{m-1}$ одержимо правило

$$P(A_1 A_2 A_3 \dots A_m) = P(A_1)P(A_2 | A_1)P(A_3 | A_1 A_2) \dots P(A_m | A_1 A_2 A_3 \dots A_{m-1}), \quad (3.16)$$

яке формулюється в такий спосіб: **ймовірність спільного настання декількох подій дорівнює добутку умовних ймовірностей цих подій, причому умовна ймовірність кожної наступної події обчислюється з урахуванням настання попередніх подій.**

Приклад 3.12. (Вентцель, 1970). Із коробки, що містить 4 білі і 3 чорні кулі, виймають (одночасно або послідовно) дві кулі. Знайти ймовірність того, що обидві кулі будуть білими.

Подамо подію $C = \{ \text{обидві кулі білі} \}$ як добуток двох подій $C = C_1 C_2$, де

$$C_1 = \{ \text{перша куля біла} \},$$

$$C_2 = \{ \text{друга куля біла} \}.$$

Тоді $P(C) = P(C_1 C_2) = P(C_1) \cdot P(C_2 | C_1)$. Очевидно, що $P(C_1) = 4/7$. Знайдемо $P(C_2 | C_1)$. Для цього припустимо, що подія C_1 вже відбулася, тобто перша куля була біла. Після цього в коробці залишилося 6 куль, із яких 3 білі, отже $P(C_2 | C_1) = 3/6 = 1/2$. Звідси $P(C) = (4/7)(1/2) = 2/7$. Той же результат отриманий у прикладі 3.9 шляхом безпосереднього підрахунку кількості випадків. Є сенс порівняти два способи розв'язання. ●

Приклад 3.13. (Вентцель, 1970). У коробці 4 білі кулі й 3 чорні. Із неї виймають по одній дві кулі. Знайти ймовірність того, що вони будуть різних кольорів.

Подія $D = \{ \text{кулі різних кольорів} \}$ розпадається на суму двох несумісних варіантів D_1 і D_2 :

$$D_1 = \{ \text{перша куля біла, друга куля чорна} \};$$

$$D_2 = \{ \text{перша куля чорна, друга куля біла} \}.$$

Ймовірність кожного з варіантів знайдемо за правилом множення. Ймовірність події D_1 дорівнює добутку ймовірності того, що перша куля біла, та умовної ймовірності того, що друга куля чорна, за умови, що перша – біла:

$$P(D_1) = (4/7) \cdot (3/6) = 2/7.$$

Так само обчислюється й ймовірність другого варіанта:

$$P(D_2) = (3/7) \cdot (4/6) = 2/7.$$

За правилом додавання ймовірностей $P(D) = P(D_1) + P(D_2) = 4/7$. ●

Приклад 3.14. (Вентцель, 1970). У коробці 5 перенумерованих куль із номерами 1, 2, 3, 4, 5. Із коробки одна за однією виймаються всі 5 куль. Знайти ймовірність того, що їх номери будуть іти в порядку зростання.

Нехай подія A_i полягає в тому, що i -та вийнята куля має номер i . Тоді нас цікавить подія $A = A_1 A_2 A_3 A_4 A_5$. За правилом добутку ймовірностей для декількох подій

$$P(A) = (1/5) \cdot (1/4) \cdot (1/3) \cdot (1/2) \cdot (1/1) = 1/120. \bullet$$

Приклад 3.15. Для оцінки ймовірності існування у Всесвіті розумних істот застосовується формула Дрейка⁶:

$$\frac{m_{\oplus}}{m} = P(A) \cdot P(B|A) \cdot P(C|AB) \cdot P(D|ABC) \cdot \frac{t}{T},$$

де $m_{\oplus}$ – кількість зірок, які, ймовірно, мають планети з одночасно існуючими високорозвиненими цивілізаціями;

m – загальна кількість зірок у Всесвіті;

A – подія, яка полягає в тому, що зірка має планетну систему;

$B|A$ – подія, яка полягає в тому, що в планетній системі виникає життя;

$C|AB$ – подія, що полягає в тому, що в процесі еволюції в планетній системі виникає цивілізація;

$D|ABC$ – подія, яка полягає в тому, що в процесі розвитку розумні істоти створюють високорозвинену цивілізацію;

t – час життя цивілізацій;

T – вік Всесвіту.

Сучасні оцінки для ймовірностей подій, що входять у формулу Дрейка, такі: $P(A)=0,01 - 0,1$; $P(B|A)=0,1$; $P(C|AB) \approx 1$; $P(D|ABC) \approx 1$, $t/T=0,5$. Ймовірність події A оцінювалася за астрономічними даними й існуючими космогонічними теоріями (зверніть увагу, що ймовірність події A визначена в спосіб, який не відрізняється від статистичного, класичного й геометричного). Ймовірність події $B|A$ визначена як середня величина оцінок, даних великою кількістю опитаних фахівців (взагалі кажучи, це не ймовірність). Оцінка подій $C|AB$ і $D|ABC$ як вірогідних пояснюється тим, що і поява розуму, і розвиток науки, техніки, технології, культури й інших атрибутів високорозвиненої цивілізації узгоджуються з дією основних механізмів еволюції, завжди спрямованих на збереження виду. При оцінці відношення t/T (це, очевидно, геометрична ймовірність) оптимістична думка полягала в тому, що час життя цивілізацій у районі певної зірки обмежується лише її “строком життя” і, як правило, становить кілька мільярдів років. Виходячи з цього, вважають, що половина часу йде на створення цивілізації.

⁶ Френсіс Дрейк – американський астроном, який здійснив першу програму пошуку радіосигналів від позаземних цивілізацій SETI – Searching for ExtraTerrestrial Intelligence в 1965 – 1969 рр.

З урахуванням усіх оцінок відносна частота зірок, які мають планети з одночасно існуючими високорозвиненими цивілізаціями, становить

$$\frac{m_{\oplus}}{m} = (0,01 - 0,1) \cdot 0,1 \cdot 1 \cdot 0,5 = 5 \cdot 10^{-2} - 5 \cdot 10^{-3}.$$

За сучасними даними, в нашій Галактиці $m \approx 10^{11}$ зірок, тоді $m_{\oplus} \approx 5 \cdot 10^8 - 5 \cdot 10^9$; у Всесвіті відповідно $m \approx 10^{21}$; $m_{\oplus} \approx 5 \cdot 10^{18} - 5 \cdot 10^{19}$. ●

3.16. Незалежні події

Перш ніж увести нове поняття незалежних подій, з'ясуємо сутність самого поняття “незалежність”. Розглянемо два довільних процеси A і B (явища, феномени й т. ін.).

Будемо вважати, що **процес A не залежить** від процесу B , якщо в проходженні процесу A при “увімкненому” й “вимкнутому” процесі B немає розбіжностей. Варто мати на увазі, що “увімкнення” або “вимикання” процесу B не слід розуміти буквально. У дослідника для цього може не бути фізичної можливості. Мається на увазі те, що спостереження за перебігом процесу A потрібно проводити в ті періоди, коли процес B через природні причини відбувається або не відбувається.

Єдиною кількісною характеристикою випадкової події як процесу є ймовірність, тому для формулювання відсутності впливу випадкової події B на подію A візьмемо таку рівність:

$$P(A|B) = P(A|\bar{B}). \quad (3.17)$$

Цю рівність можна вважати **першим визначенням незалежності випадкових подій**: випадкова подія A не залежить від випадкової події B , якщо умовна ймовірність події A при настанні події B дорівнює умовній ймовірності події A при ненастанні події B .

Одержимо інше формулювання незалежності випадкових подій. Для цього розглянемо розбиття достовірної події за допомогою прямої та оберненої подій:

$$\Omega = B + \bar{B}.$$

Як розбиття універсальної множини породжує розбиття будь-якої її підмножини (див. розд. 1.3), так і розбиття вірогідної події породжує розбиття будь-якої можливої в умовах даного досліду події. Застосуємо дану властивість для розбиття події A :

$$A = A\Omega = A(B + \bar{B}) = AB + A\bar{B}.$$

Оскільки AB і $A\bar{B}$ – несумісні події, тобто варіанти події A , то можна застосувати правило додавання ймовірностей для двох несумісних подій:

$$P(A) = P(AB) + P(A\bar{B}).$$

Звідси

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

Остання рівність потрібна для обчислення умовних ймовірностей у визначенні незалежності (3.17). Дійсно, для лівої й правої частин (3.17) з урахуванням (3.15) відповідно одержуємо

$$\frac{P(AB)}{P(B)} = \frac{P(A\bar{B})}{P(\bar{B})},$$

$$\frac{P(AB)}{P(B)} = \frac{P(A) - P(AB)}{1 - P(B)}.$$

В останньому виразі в знаменнику правої частини застосований перший висновок із правила додавання ймовірностей для несумісних подій. Після зведення останнього виразу до спільного знаменника отримаємо

$$P(AB) - P(AB)P(B) = P(A)P(B) - P(AB)P(B),$$

звідки після скорочення виведемо співвідношення, яке застосуємо для нового визначення незалежності випадкових подій:

$$P(AB) = P(A)P(B).$$

Аналогічно може бути отримане ще одне співвідношення:

$$P(A\bar{B}) = P(A)P(\bar{B}).$$

За допомогою цих співвідношень та з урахуванням (3.15) можна дати **друге визначення незалежності випадкових подій**:

$$P(A|B) = P(A|\bar{B}) = P(A), \quad (3.18)$$

яке формулюється так: випадкова подія A не залежить від випадкової події B , якщо умовна ймовірність події A при настанні або ненастанні події B дорівнює безумовній ймовірності події A .

3.17. Правило множення ймовірностей для двох незалежних подій

Одержана при формулюванні другого визначення незалежності випадкових подій допоміжна рівність іноді використовується як самостійна властивість незалежних подій за назвою правила множення: **ймовірність спільного настання двох незалежних подій дорівнює добутку їх ймовірностей**:

$$P(AB) = P(A)P(B). \quad (3.19)$$

Приклад 3.16. (Вентцель, 1970). Дослід полягає в підкиданні двох монет. Перевірити незалежність подій

$A = \{ \text{поява "орла" на першій монеті} \},$

$B = \{ \text{поява "орла" на другій монеті} \}.$

Простором елементарних подій для даного досліду є $\Omega = \{OO, OP, PO, PP\}$, події, що розглядаються: $A = \{OO, OP\}$, $B = \{OO, PO\}$, а їх добуток $AB = \{OO\}$. За “класичною” формулою $P(A) = 1/2$, $P(B) = 1/2$, $P(AB) = 1/4$, і

$$P(A|B) = \frac{P(AB)}{P(B)} = 1/2 = P(A),$$

$$P(B|A) = \frac{P(AB)}{P(A)} = 1/2 = P(B)$$

Отже, події A і B незалежні. ●

Приклад 3.17. (Гільдерман, 1991). У групі обстежених 1 000 чоловік. Із них 600 палять і 400 не палять. Серед тих, хто палить, 240 чоловік мають ті або інші захворювання легенів. Серед тих, хто не палить, захворювання легенів мають 120 чоловік. Чи є паління й захворювання легенів незалежними подіями?

Розглянемо події

$A = \{ \text{обстежений палить} \},$

$B = \{ \text{обстежений має захворювання легенів} \}.$

Тоді $P(AB) = 240/1000 = 0,24$; $P(A) = 600/1000 = 0,6$; $P(B) = (240 + 120)/1000 = 0,36$. Оскільки $P(AB) = 0,24 \neq P(A)P(B) = 0,6 \cdot 0,36$, то події A і B залежні. Варто звернути увагу на таке. Якщо обмежити відповідь на питання тільки групою обстежених, то розглянуті ймовірності, зрозуміло, є класичні. Якщо ж відповідь поширюється на людей і паління взагалі, то розглянуті ймовірності є відносними частотами, причому про стабілізацію нічого не відомо. ●

3.18. Правило множення ймовірностей для декількох незалежних подій

Із правила множення для двох незалежних подій методом індукції легко виводиться узагальнення для декількох незалежних подій $A_1, A_2, A_3, \dots, A_m$: **ймовірність добутку декількох незалежних подій дорівнює добутку їх ймовірностей**:

$$P(A_1 A_2 \dots A_m) = P(A_1)P(A_2) \dots P(A_m). \quad (3.20)$$

3.19. Двохетапні дослід

Розглянемо дослід, що проходить у два етапи.

На **першому етапі** настає одна з подій $H_1, H_2, \dots, H_m$, які називаються **гіпотезами** і для яких виконуються такі умови:

- а) гіпотези утворюють повну групу $H_1 + H_2 + \dots + H_m = \Omega$;
- б) гіпотези несумісні, тобто $H_i H_j = \emptyset$, $i, j = \overline{1, m}$, $i \neq j$;
- в) відомі ймовірності гіпотез $P(H_1), P(H_2), \dots, P(H_m)$.

На **другому етапі**, після настання однієї з гіпотез, може статися (а може й не статися) деяка подія A , причому є відомі умовні ймовірності настання цієї події після настання гіпотез:

$$P(A|H_1), P(A|H_2), \dots, P(A|H_m).$$

Схема двухетапного дослід показана на рис. 3.5 і 3.6.

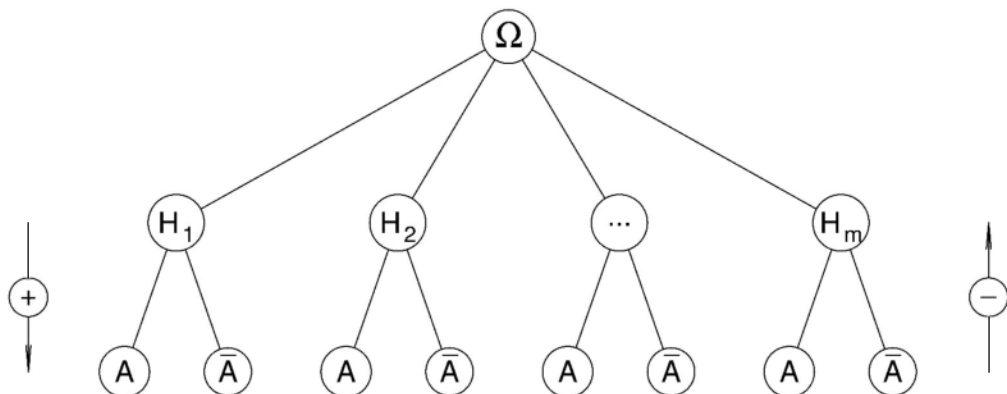
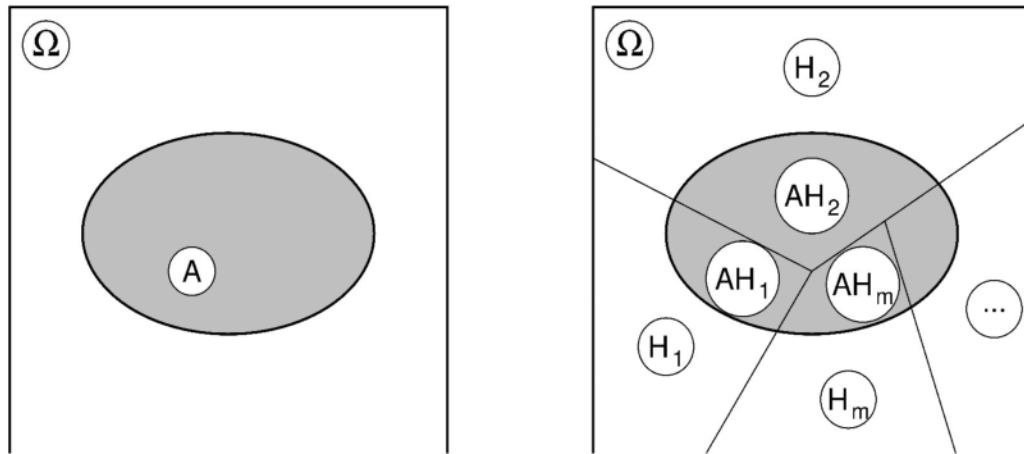


Рис. 3.5. Схема двухетапного дослід (настання гіпотез H_i передую настанню події A або події $\bar{A}$).

Рис. 3.6. Розбиття події A на варіанти за допомогою гіпотез

У двохетапних дослідках звичайно розглядають два типи задач, залежно від обраного напрямку. При прямому ході, умовно позначеному на рис. 3.5 стрілкою з “плюсом”, виходячи з відомостей про двохетапний дослід, потрібно визначити безумовну ймовірність настання події A (апостеріорну або переддослідну), тоді як при зворотньому ході, позначеному на рис. 3.5 стрілкою з “мінусом”, потрібно на основі факту настання події A визначити умовні ймовірності гіпотез (апостеріорні, або післядослідні). Ці задачі розв’язуються за допомогою формули повної ймовірності і формули Байєса. Подамо вивід цих формул одночасно: формули повної ймовірності – ліворуч, формули Байєса – праворуч.

Оскільки гіпотези розбивають вірогідність події Ω , то автоматично буде здійснене й розбиття події A :

$$A = A\Omega = A(H_1 + H_2 + \dots + H_m) = AH_1 + AH_2 + \dots + AH_m$$

Застосувавши правило додавання ймовірностей для декількох несумісних подій (3.12) обчислимо ймовірність події A :

$$P(A) = P(AH_1) + P(AH_2) + \dots + P(AH_m).$$

Для обчислення ймовірностей варіантів AH_i події A застосуємо правило множення ймовірностей для двох довільних подій (3.15):

Вихідним пунктом, як і при виведенні формули повної ймовірності, є розбиття простору елементарних подій Ω за допомогою гіпотез:

$$\Omega = H_1 + H_2 + \dots + H_m.$$

Розглянемо один із варіантів AH_i події A і обчислимо його ймовірність $P(AH_i)$. Оскільки події A і H_i взагалі є сумісні (рис. 3.6), то варто застосувати правило множення ймовірностей для двох довільних подій (3.15):

$$P(AH_i) = P(A)P(H_i | A) = P(H_i)P(A | H_i).$$

Після ділення обох частин останньої рівності на $P(A)$ одержуємо шукану

$$P(AH_1) = P(A|H_1)P(H_1),$$

$$P(AH_2) = P(A|H_2)P(H_2),$$

$$P(AH_m) = P(A|H_m)P(H_m).$$

Підставивши останні рівності у вираз для ймовірності події A , одержуємо формулу повної ймовірності:

$$P(A) = P(A|H_1)P(H_1) + P(A|H_2)P(H_2) + \dots + P(A|H_m)P(H_m)$$

або

$$P(A) = \sum_{i=1}^m P(H_i)P(A|H_i) \quad (3.21).$$

умовну ймовірність гіпотези H_i при настанні події A :

$$P(H_i|A) = \frac{P(H_i)P(A|H_i)}{P(A)}. \quad (3.21)$$

Останній вираз і є формулою Байєса, однак її можна застосовувати тільки за умови, що знаменник обчислюється за формулою повної ймовірності (3.20). Крім того, потрібно пам'ятати, що формула Байєса застосовується для кожної гіпотези, тобто індекс i набуває значень від 1 до m .

Приклад 3.18. (Гільдерман, 1991). У трьох однакових коробках містяться: у першій – 2 білі й 3 чорні кулі, у другій – 4 білі й 1 чорна, у третій – 3 білі. З однієї коробки навмання виймається одна куля. Знайти ймовірність того, що ця куля буде білою, тобто ймовірність події

$$A = \{ \text{вийнята куля біла} \}.$$

Можливі три гіпотези:

$$H_1 = \{ \text{обрана перша коробка} \};$$

$$H_2 = \{ \text{обрана друга коробка} \};$$

$$H_3 = \{ \text{обрана третя коробка} \}.$$

За умовою задані ймовірності гіпотез $P(H_1) = P(H_2) = P(H_3) = 1/3$ і умовні ймовірності $P(A|H_1) = 2/5$, $P(A|H_2) = 4/5$, $P(A|H_3) = 1$. За формулою повної ймовірності (3.20):

$$P(A) = P(H_1)P(A|H_1) + P(H_2)P(A|H_2) + P(H_3)P(A|H_3) =$$

$$= (1/3) \cdot (2/5) + (1/3) \cdot (4/5) + (1/3) \cdot 1 = 11/15. \quad \bullet$$

Приклад 3.19. (Гільдерман, 1991). У трьох однакових коробках містяться: у першій – 3 білі й 1 чорна куля; у другій – 2 білі й 3 чорні; у третій – 3 білі. З однієї коробки навмання виймається 1 куля. Ця куля виявилася білою. Знайти апостеріорні ймовірності того, що куля вийнята з першої, другої або третьої коробки.

Оскільки коробка вибирається навмання, то апріорні ймовірності гіпотез рівні:

$$P(H_1) = P(H_2) = P(H_3) = 1/3,$$

а умовні ймовірності події

$$A = \{\text{вийнята біла куля}\}$$

при гіпотезах H_1, H_2, H_3 такі:

$$P(A|H_1) = 3/4; \quad P(A|H_2) = 2/5; \quad P(A|H_3) = 1.$$

За формулою Байєса (3.21) апостеріорні ймовірності гіпотез H_1, H_2, H_3 дорівнюють:

$$P(H_1|A) = \frac{(1/3) \cdot (3/4)}{(1/3) \cdot (3/4) + (1/3)(2/5) + (1/3) \cdot 1} = \frac{15}{43},$$

$$P(H_2|A) = \frac{8}{43},$$

$$P(H_3|A) = \frac{20}{43}.$$

Оскільки гіпотези несумісні й утворюють повну групу подій (див. розд. 3.4), то умовну ймовірність $P(H_3|A)$ можна обчислити простіше:

$$P(H_3|A) = 1 - P(H_1|A) - P(H_2|A) = 1 - 15/43 - 8/43 = 20/43.$$

З урахуванням результату дослідів (настання події A) ймовірності гіпотез змінилися: найбільш ймовірною стала гіпотеза H_3 , найменш ймовірною – гіпотеза H_2 . ●

4. Випадкові величини

4.1. Поняття випадкової величини

Уведене в розд. 3 поняття випадкової події не вичерпує всіх можливих ситуацій у досліді із випадковим результатом. Нагадаємо, що в подібних досліді спостереження за їх результатом рівнозначне розпізнаванню двох якісно різних ситуацій: події A та протилежної їй події $\bar{A}$, тобто наявності результату, який становить певний інтерес, або його відсутності. Але для дослідів, у яких результат є наслідком виміру деякої величини або показника досліджуваного процесу, випадковість виявляється як непередбачуваність, невизначеність результату виміру в серії дослідів із незмінними умовами їх проведення. Звернемо увагу на те, що подібні досліді **завжди** закінчуються деяким **кількісним результатом** – результатом вимірювання. Це й відрізняє їх від дослідів, у яких спостерігається **якісний результат**, тобто випадкові події.

Стосовно дослідів із кількісним результатом відзначимо таке.

По-перше, як і в досліді із випадковими подіями, повинен бути деякий прихований механізм, що породжує відгук, який реєструється (результат виміру), на певне збурення умов досліді.

По-друге, відгук, тобто кількісний результат досліді, не може бути довільним, тому що для будь-якого досліджуваного процесу його кількісні показники завжди мають мінімальне й максимальне значення. Отже, можна говорити про певну множину допустимих значень результатів досліді.

Звідси випливає таке визначення випадкової величини.

Під **випадковою величиною** в умовах даного досліді будемо розуміти деяку множину значень і прихований механізм, що в сукупності приводять до появи (спостереження) одного з можливих значень, як відгук на збурення, в окремому досліді (рис. 4.1). За цих умов цілком зберігається визначення досліді, наведене в розд. 3.1.

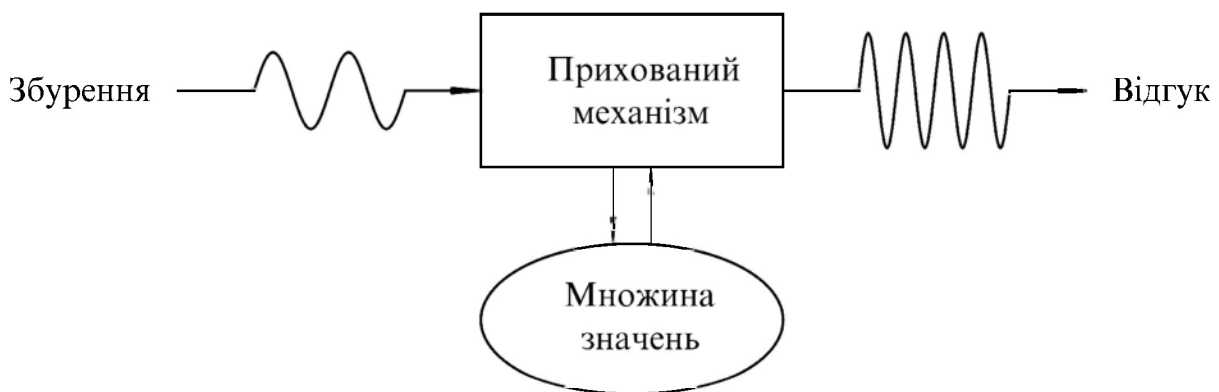


Рис. 4.1. Схема до визначення випадкової величини

Подібні визначення вводяться, як правило, тоді, коли структура досліджуваної системи або невідома (прихована), або складна (якщо взагалі можлива) для **детермінованого аналізу** (тобто такого, у якому вивчаються механізми типу “певний вплив – певний відгук”). У таких ситуаціях говорять, що система (процес) заміню-

ється моделлю “чорного ящика”, а єдиним способом вивчення системи стає **стохастичний аналіз** (спостереження за її відгуками на задані впливи й спроба пошуку закономірностей у відгуках).

4.2. Позначення, пов'язані з випадковими величинами

Під час роботи з випадковими величинами будемо використовувати два способи позначень (табл. 4.1). Залежно від конкретної ситуації певні переваги має або перший, або другий спосіб.

Таблиця 4.1

Спосіб 1	Спосіб 2
1. Великі літери $X, Y, \dots$ використовуються для випадкових величин (запасу значень і прихованого механізму). 2. Малі літери використовуються для реєстрованих у досліді значень цих величин: $x_1, x_2, \dots, x_{100}, \dots;$ $y_1, y_2, \dots, y_{100}, \dots;$ де 1, 2, ... – номер спостереження.	1. Малі літери $x_1, x_2, \dots$ – для значень, що спостерігаються в досліді (тобто це такі самі позначення, як і в п. 2 способу 1). 2. Малі літери (найчастіше з індексами) з тильдою ($\sim$) зверху: $\tilde{x}_1, \tilde{x}_2, \dots$ – для випадкових величин.

4.3. Класифікація випадкових величин

Випадкові величини поділяються на дискретні й неперервні (рис. 4.2).

Дискретна випадкова величина набуває множини значень, кожне з яких відділене від інших інтервалами значень, що є неможливі.

Приклад 4.1. Нехай у серії дослідів підраховується кількість пелюсток на квітках рослин певного виду, що ростуть на деякому ареалі. Визначимо випадкову величину X , значення якої в одному досліді дорівнює кількості пелюсток для конкретної рослини. Зрозуміло, що кількість пелюсток не менша одиниці і не більша деякого максимального значення, але вона не може дорівнювати, наприклад, 3,5 або 4,87. Тому інтервали значень, розташованих між цілими числами, тобто вигляду (2,3), (5,6) і т.д., є принципово неможливі для випадкової величини X . ●

Зауваження 4.1. Із визначення дискретної випадкової величини не впливає, що множина її значень утворена тільки цілими числами, хоча такі дискретні випадкові величини, мабуть, найчастіше й зустрічаються. Приклад 4.1 не повинен сприйматися як типовий. Випадкова величина може набувати не тільки цілих значень, однак вони відділені одне від одного інтервалами неможливих значень. ▲



Рис. 4.2. Види випадкових величин

Неперервна випадкова величина набуває будь-яких значень із деякого скінченного або нескінченного інтервалу. Мається на увазі, що для будь-яких двох значень неперервної випадкової величини x_1 і x_2 , розташованих як завгодно близько, можна вказати третє значення випадкової величини x_3 , розташоване між ними (рис. 4.3). Для цього проміжного значення та одного з двох крайніх – ліворуч або праворуч від нього (на рис. 4.3 обрана пара значень x_2 і x_3) – знову можна вказати значення випадкової величини x_4 , розташоване між ними, і т.д.

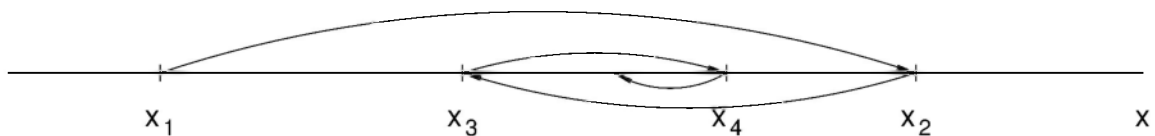


Рис. 4.3. Схема до визначення неперервної випадкової величини

Зауваження 4.2. Неперервна випадкова величина є абстракцією, ідеалізацією, моделлю, нехай навіть досить правдоподібною й конструктивною (тобто такою, яка дозволяє робити практично важливі висновки). У науці часто застосовуються моделі реального об'єкта, що зберігають його основні властивості й нехтують другорядними, наприклад: ідеальний газ, абсолютно тверде тіло, абсолютно пружне тіло, нерозтяжна нитка, рух без тертя і т.д. У моделі неперервної випадкової величини також збережені основні властивості реального об'єкта. Це в першу чергу можливість щільно (без проміжків) заповнювати своїми значеннями деякий інтервал. Точніше, мова йде про відсутність формального обмеження на набуття будь-яких значень із деякого інтервалу. Насправді в реального об'єкта така властивість у “чистому” вигляді може й не спостерігатися, наприклад, через неможливість проводити необмежену кількість дослідів протягом необмеженого часу. Крім того, сам процес виміру здійснюється з обмеженою точністю, тобто зі скінченною кількістю знаків. Тому в інтервалі можливих значень залишаються невеликі “просвіти”. Однак якщо відмовитися від моделі неперервної випадкової величини, то виявиться, що чимало реальних процесів є досить складними для вивчення, оскільки унеможлиблюється застосування розвиненого математичного апарата (функції, границі, похідні й т.ін). ▲

У досліді неперервна випадкова величина спостерігається як **інтервальна випадкова величина**. Для інтервальної випадкової величини властиве таке:

1) у множині можливих значень, розділених на інтервали (такі інтервали можуть бути утворені шкалою вимірювального приладу, але не обов'язково нею), можна вказати інтервал, у якому знаходиться результат даного досліду;

2) у межах даного інтервалу результат виміру не локалізується (усі значення, що потрапляють у даний інтервал, не розрізняються).

Приклад 4.2. Нехай у досліді вимірюється зріст людей. Запас значень відповідної випадкової величини X визначається групою людей, у межах якої проводиться дослід. Вимірний зріст однієї людини – значення x , набуте випадковою величиною X в окремому досліді. Якщо використовується найпростіший вимірювальний прилад – лінійка із сантиметровими поділками, то інтервали, про які йде мова у визначенні інтервальної випадкової величини, такі: $[159,5 \text{ см}, 160,5 \text{ см})$, $[160,5 \text{ см}, 161,5 \text{ см})$, $[161,5 \text{ см}, 162,5 \text{ см})$. Дійсно, нехай результат виміру зросту деякої людини знаходиться між двома поділками лінійки – 159 см і 160 см. На око можна досить точно визначити, до якої з цих двох поділок ближчий результат (значення зросту). Відповідно до відомого правила округлювання результат виміру, ближчий до лівої межі ($x < 159,5 \text{ см}$), округлюється з нестачею, тобто до 159 см, а ближчий до правої межі ($x \geq 159,5 \text{ см}$) – з надлишком, тобто до 160 см. Так само округлюються результати виміру, що знаходяться між 160 см і 161 см. Тому при роботі з даним вимірювальним приладом зріст 160 см приписується всім людям, дійсний зріст яких виявився в інтервалі $[160,5 \text{ см}, 161,5 \text{ см})$. Оскільки десяті, соті та інші частини однієї поділки, яка дорівнює 1 см, не визначаються, то всі результати вимірів у межах даного інтервалу стають зовсім рівноправними. Щоправда, можна сказати, більші вони чи менші 160 см, тобто чи знаходяться праворуч або ліворуч від цієї поділки, але це з погляду визначення інтервальної випадкової величини не має значення. ●

4.4. Опис випадкової величини

Оскільки введення поняття випадкової величини (моделі “чорної ящика”) спрямоване на вивчення деякої системи (процесу), варто відразу вказати кінцеву мету цього вивчення. Нею є опис випадкової величини. В описі випадкової величини є деяка подібність з описом випадкової події. Саме необхідність опису випадкової події A веде до введення поняття її ймовірності $P(A)$ (див. розд. 3.3). При статистичному трактуванні ймовірність $P(A)$ встановлює закономірність спостережень події A в деякій серії дослідів. Зрозуміло, що численність результатів спостереження випадкової величини (тобто спостереження не двох протилежних станів, як у випадку спостереження випадкової події, тобто A і $\bar{A}$, а наявність саме множини значень) – свідчення того, що її опис не може зводитися до одного числа, а певним чином залежить від множини (запасу) значень. А оскільки множини значень дискретної випадкової величини й неперервної випадкової величини влаштовані по-різному (одна множина дискретна, тобто з розривами, а друга – неперервна, тобто суцільна), то мають розрізнятися описи дискретних і неперервних випадкових величин.

Повний опис випадкової величини – це деяке правило, що дозволяє обчислювати ймовірності набуття дискретною випадковою величиною певних значень і неперервною випадковою величиною – значень із певного інтервалу. Повний опис ви-

падкової величини конкретизується загальним поняттям **функції розподілу** (для дискретної і неперервної випадкових величин, див. розд. 4.11) і двома частковими поняттями : **ряду розподілу** – для дискретної випадкової величини і **функцією густини** – для неперервної випадкової величини. Зазначимо, що у разі повного опису випадкової величини використовують ще термін **закон розподілу**.

На противагу повному опису досить часто розглядають **неповний опис випадкової величини**. Він вводиться за допомогою **числових характеристик випадкової величини**.

4.5. Поняття генеральної та вибіркової сукупностей

У дослідях множина значень випадкової величини (не обов'язково повністю) спостерігається, як правило, у вигляді числових значень деякої кількісної ознаки, носіями якої є об'єкти, що, у свою чергу, належать до деякої множини. Ця множина називається **генеральною сукупністю**. Чисельність генеральної сукупності позначається великою літерою N і називається **обсягом генеральної сукупності**. Для вивчення випадкової величини (деякої ознаки) із генеральної сукупності вибирається деяка підмножина, що називається **вибірковою сукупністю** або **вибіркою**. Чисельність вибіркової сукупності позначають малою літерою n і називають **обсягом вибіркової сукупності**. Вибіркові значення будемо позначати в такий спосіб:

$$x_{i=1} = x_1, x_{i=2} = x_2, \dots x_i, \dots x_{i=n} = x_n \text{ або } x_i, i = \overline{1, n}.$$

Зауваження 4.3. Нижній індекс i при малій позначці випадкової величини (згадайте прийняту систему позначень випадкових величин, наведену в розд. 4.2) буде застосовуватися **тільки для позначення вибірових значень**. ▲

Поняття вибіркової та генеральної сукупностей нерозривні. Називаючи генеральну сукупність (об'єкт вивчення), слід указати вибірку сукупність (засіб вивчення генеральної сукупності); називаючи вибірку сукупність, потрібно вказати, із якої генеральної сукупності вона була витягнута, оскільки висновки, що робляться за результатами вивчення вибірки, поширюються на всю генеральну сукупність.

4.6. Класифікація вибірок (методи витягнення вибірок)

Прийнято розрізняти вибірки (Іванова, 1987) (див. рис. 4.4) за:

- 1) обсягом;
- 2) способом відбору;
- 3) схемою випробувань.

Малі вибірки використовуються для статистичного контролю вже відомих властивостей генеральної сукупності.

Великі вибірки використовуються для виявлення невідомих властивостей генеральної сукупності.

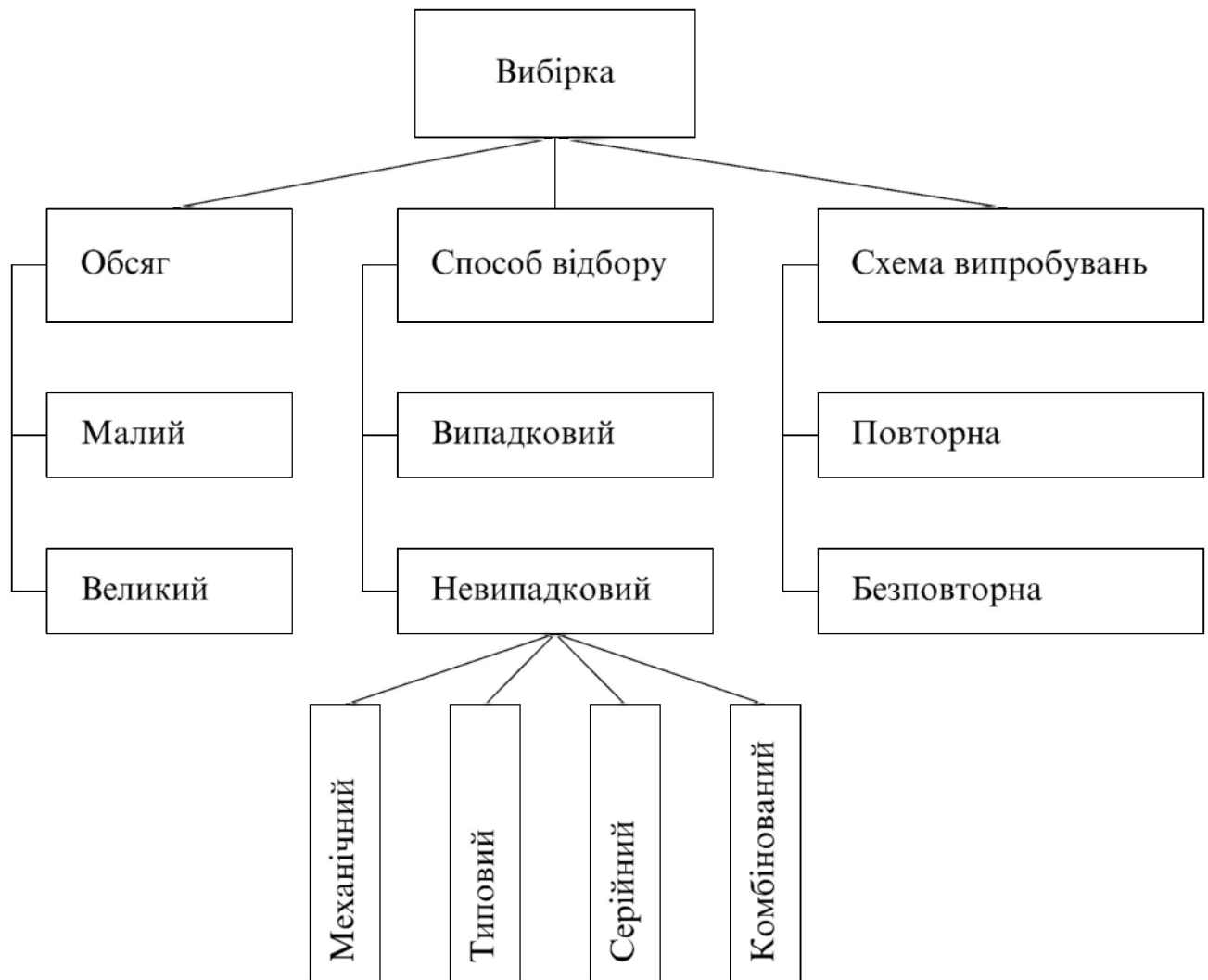


Рис. 4.4. Схема класифікації вибірок

Випадковий спосіб відбору полягає в тому, що створюються умови, за яких усі об'єкти генеральної сукупності мають однакову ймовірність (однакові шанси) бути включеними у вибірку (це ідеальна, а тому недосяжна модель вибору, оскільки генеральна сукупність повинна бути схожа на коробку з добре перемішаними кульками, а вибіркoва сукупність – на деяке сполучення з N елементів-кульок генеральної сукупності по n елементів-кульок).

Невипадковий спосіб відбору означає будь-яке наближення до випадкового способу добору, що може бути здійснений.

Механічний спосіб відбору (рис. 4.5) полягає в тому, що генеральна сукупність поділяється на деяку кількість приблизно рівновеликих частин, із яких у вибірку витягається “випадково” однакова кількість об'єктів. При поділі генеральної сукупності на нерівновеликі частини у вибіркoву сукупність витягається кількість об'єктів, пропорційна чисельності даної частини.

Типовий спосіб відбору (рис. 4.5) полягає в тім, що генеральна сукупність поділяється на деяку кількість однорідних (гомогенних) за деякою ознакою (що не вивчається у даному досліді) частин, чисельності яких загалом неоднакові, після чого

з кожної гомогенної частини “випадково” витягається у вибірку кількість об'єктів, пропорційна чисельності цієї частини.

Серійний спосіб відбору (рис. 4.5) полягає в тому, що генеральна сукупність поділяється на певну кількість груп, названих серіями, і деякі з них повністю включаються у вибірку.

Комбінований спосіб відбору являє собою можливе й доцільне в даному досліді сполучення трьох вищезазначених способів добору.

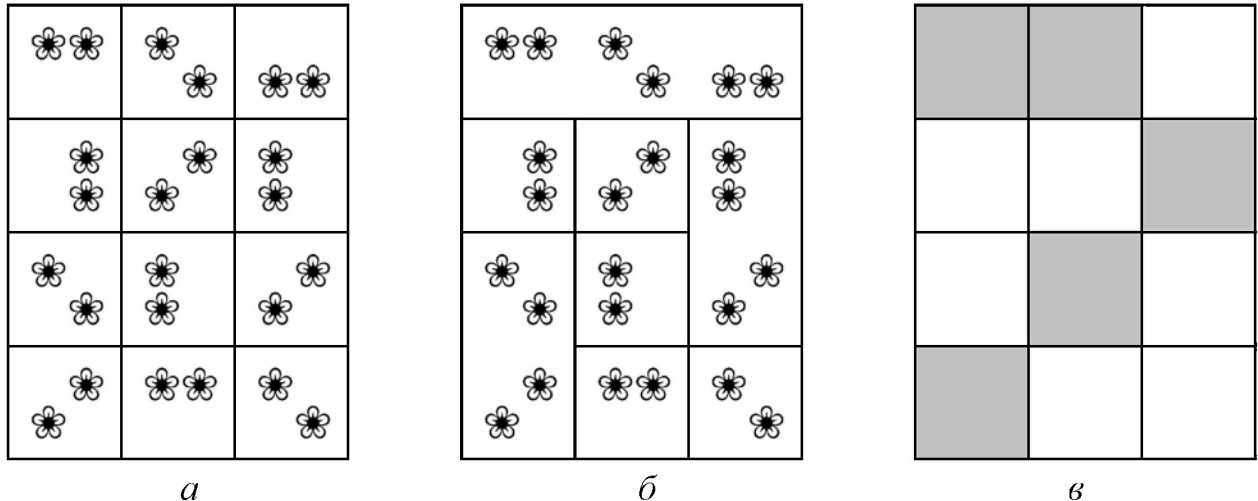


Рис. 4.5. Невипадкові способи вибору: *a* – механічний, *б* – типовий, *в* – серійний

Схема називається **повторною**, якщо об'єкт генеральної сукупності після витягнення у вибірку повертається в генеральну сукупність і може бути знову витягнутий при наступних дослідях.

Схема випробувань називається **неповторною**, якщо об'єкт генеральної сукупності після включення у вибірку не повертається в генеральну сукупність.

При великому обсязі генеральної сукупності N , точніше, при малій величині відношення n/N , розходження між повторною і безповторною схемами випробувань практично відсутнє.

Вибірка називається **репрезентативною (показовою)**, якщо вона в кількісному і якісному відношеннях правильно відображає генеральну сукупність. У противному разі відбувається процес **зсуву**. У результаті вибірка представляє не всю генеральну сукупність, а лише якийсь її шар, тобто лише її частину. Зсув – одне з основних джерел помилок при застосуванні вибіркового методу.

Приклад 4.3 (Гласс, 1976). У 1936 р. кандидатами на президентських виборах у США були Ф.Д. Рузвельт і А.М. Ландон. Напередодні виборів американський журнал “Літературний огляд” провів опитування щодо результату цих виборів. Використавши телефонні книги, редакція відібрала випадково 4 мільйони (!) адрес, на які вона розіслала по всій країні листівки з питаннями про ставлення до кандидатів у президенти. Затративши велику суму грошей на розсилку й обробку листівок, журнал оголосив, що на майбутніх виборах президентом США з великою перевагою буде обраний А.М. Ландон. Результат виборів виявився протилежним цьому прогнозові.

У даному випадку були допущені відразу дві помилки. По-перше, телефонні

книги не давали репрезентативну вибірку з населення країни, хоч би тому що абоне-нтами були в основному заможні родини. По-друге, надіслали відповіді не всі, а лю-ди, досить упевнені у своїй думці, звиклі відповідати на листи, тобто здебільшого представники ділового світу, які й підтримували А.М. Ландона.

Водночас соціологи Дж. Геллап і Е. Роупер правильно пророкували перемогу Ф.Д. Рузвельта, ґрунтуючись тільки на 4 тисячах анкет. Запорукою цього успіху, що зробив славу його авторам, було не тільки правильне складання вибірки. Вони вра-хували, що суспільство розпадається на соціальні групи, більш-менш однорідні сто-совно політичних уподобань. Тому навіть нечисленна вибірка з шару може бути по-казовою порівняно з численною вибіркою без урахування розшарування генеральної сукупності. Маючи результати обстеження по шарах, можна характеризувати суспі-льство в цілому. Зараз така методика є загальноприйнята. ●

Приклад 4.4 (Глотов, 1985). При вивченні віддалених наслідків черепно-мозкових травм у базі даних травматологічного інституту були відібрані історії хво-роб осіб, які потрапили в автомобільну катастрофу 10 – 15 років тому. Контрольна група була представлена людьми, які одержали якусь травму, наприклад перелом кі-нцівки, при цьому записи в історії хвороби повністю заперечували можливість чере-пно-мозкової травми. “Дослідна” група, навпаки, включала людей, які перенесли че-репно-мозкові травми певної категорії. Виявилось, що в групі людей, які перенесли 10 – 15 років тому черепно-мозкові травми, набагато частіше, ніж у контрольній, зу-стрічаються шизофренія, епілепсія й інші важкі хвороби. Це явно суперечило всій сукупності наявних даних про віддалені наслідки черепно-мозкових травм. Дослі-дження було повторене шляхом витягання нової вибірки з бази даних. Цього разу та сама важка патологія зустрічалася частіше в контрольній групі. Після детального ви-вчення всіх даних було виявлено, що особи, які належали до контрольної та дослід-ної груп, розрізняються не тільки за характером травми, але й за рядом ознак, що не мають, на перший погляд, прямого зв'язку з досліджуваним питанням: співвідно-шення за статтю, розподіл за віком, професією і т. ін. Таким чином, ефекти, що спо-стерігалися, були наслідком якихось неконтрольованих у досліді причин і умов.

Тому для вивчення стали відбирати не всі історії хвороб підряд, а виходячи на-самперед із певної “модельної” ситуації: автомашини врізалася в групу людей, хтось випадково одержав черепно-мозкову травму, а хтось – інший вид ушкоджень. Такий спосіб витягнення “випадкової” вибірки дав дуже чіткі результати: у групі осіб, які перенесли черепно-мозкові травми, з більшою частотою зустрічалися відносно слаб-ко виражені психоневрологічні відхилення. ●

4.7. Повний опис дискретної випадкової величини

Нехай дана деяка вибірка x_i , $i = \overline{1, n}$. Із вибірових значень виділимо групи по-вторюваних значень, для яких уведемо позначення x_k , $k = \overline{1, K}$, де k – номер групи, а K – кількість таких груп, причому $K \leq n$. Рівність $K=n$ може мати місце тільки тоді, коли повторюваних значень немає, тобто всі вибірові значення різні.

Зауваження 4.4. Велику літеру K завжди будемо використовувати для позна-чення **кількості груп** повторюваних значень дискретної випадкової величини, малу

літеру k – для позначення **номера групи**, малу позначку випадкової величини з нижнім індексом k – для **значення групи** (див. зауваження 4.3). ▲

Приклад 4.5. У досліді одержані такі результати вимірів: 0, 1, 2, 3, 2, 1, 0, 0, 1, 3, 2, 1. Обсяг даної вибірки $n = 12$. Легко помітити, що вона утворена чотирма повторюваними значеннями:

$$x_{k=1} = 0; \quad x_{k=2} = 1; \quad x_{k=3} = 2; \quad x_{k=4} = 3,$$

тобто кількість груп $K = 4$.

Для кожної групи з номером k підрахуємо її чисельність, яку будемо називати частотою значення x_k і позначати n_k . Для пар значень (x_k, n_k) складемо **вибірковий ряд розподілу частот** (табл. 4.2).

Таблиця 4.2

x_k	$x_{k=1}$	$x_{k=2}$	...	x_k	...	x_K
n_k	$n_{k=1}$	$n_{k=2}$	...	n_k	...	n_K

Зауваження 4.5. Зверніть увагу на те, що сума частот n_k у другому рядку вибіркового ряду розподілу частот (табл. 4.2) дорівнює обсягу вибірки n :

$$\sum_{k=1}^n n_k = n,$$

тобто обсяг вибірки “розмазаний”, або розподілений за значеннями x_k . Звідси й назва таблиці. ▲

Приклад 4.5 (продовження). Для розглядуваної вибірки наведемо вибірковий ряд розподілу частот (табл. 4.3).

Таблиця 4.3

x_k	0	1	2	3
n_k	3	4	3	2

Якщо від частот n_k перейти до відносних частот n_k / n , то вийде **вибірковий ряд розподілу відносних частот** (табл. 4.4).

Таблиця 4.4

x_k	$x_{k=1}$	$x_{k=2}$	...	x_k	...	x_K
$\frac{n_k}{n}$	$\frac{n_{k=1}}{n}$	$\frac{n_{k=2}}{n}$	...	$\frac{n_k}{n}$	...	$\frac{n_K}{n}$

Зауваження 4.6. Зверніть увагу на те, що сума відносних частот n_k / n у другому рядку вибіркового ряду розподілу відносних частот дорівнює 1:

$$\sum_{k=1}^K \frac{n_k}{n} = \frac{n_1}{n} + \frac{n_2}{n} + \dots + \frac{n_K}{n} = \frac{n_1 + n_2 + \dots + n_K}{n} = \frac{n}{n} = 1.$$

Ця одиниця “розмазана”, або розподілена за значеннями x_k . ▲

Приклад 4.5 (продовження). Вибірковий ряд розподілу відносних частот для розглядуваної вибірки наведений у табл. 4.5.

Таблиця 4.5

x_k	0	1	2	3
$\frac{n_k}{N}$	$\frac{3}{12}$	$\frac{4}{12}$	$\frac{3}{12}$	$\frac{2}{12}$

Якщо при збільшенні обсягу вибірки спостерігається стабілізація відносних частот n_k / n (див. розд. 3.3), то в ряду розподілу відносних частот їх можна замінити статистичними ймовірностями p_k^* (табл. 4.6).

Таблиця 4.6

x_k	$x_{k=1}$	$x_{k=2}$	...	x_k	...	x_K
p_k^*	$p_{k=1}^*$	$p_{k=2}^*$	...	p_k^*	...	p_K^*

При $n = N$ вибіркового ряду розподілу відносних частот переходить у **генеральний ряд розподілу** дискретної випадкової величини, що відрізняється від вибіркового ряду розподілу відносних частот тим, що в останньому рядку розташовуються ймовірності p_k (табл. 4.7).

Таблиця 4.7

x_k	$x_{k=1}$	$x_{k=2}$	...	x_k	...	x_K
p_k	$p_{k=1}$	$p_{k=2}$	...	p_k	...	p_K

Зауваження 4.7. Сума статистичних ймовірностей у другому рядку табл. 4.6 дорівнює одиниці:

$$\sum_{k=1}^n p_k^* = 1,$$

що випливає з подібної властивості для ряду розподілу відносних частот та визначення статистичних ймовірностей (див. розд. 3.4). Така сама властивість виконується й для ймовірностей у другому рядку табл. 4.7:

$$\sum_{k=1}^n p_k = 1.$$

Цей результат можна обґрунтувати в такий спосіб. Випадкові події $\{X = x_k\}$, $k = \overline{1, K}$ в умовах даного досліду утворюють повну групу несумісних подій, тому ймовірність суми всіх цих подій (пересвідчіться, що вони є елементарні, хоча на результаті це не позначається) дорівнює сумі їх ймовірностей. Дійсно, з умови повноти випливає, що

$$\begin{aligned} \{X = x_1\} + \{X = x_2\} + \dots + \{X = x_K\} &= \Omega, \\ P(\{X = x_1\} + \{X = x_2\} + \dots + \{X = x_K\}) &= P(\Omega) = 1, \end{aligned}$$

а умова несумісності подій дозволяє для обчислення ймовірності суми цих подій застосувати правило додавання (див. розд. 4.3), тобто

$$\begin{aligned} P(\{X = x_1\} + \{X = x_2\} + \dots + \{X = x_K\}) &= \\ = P(\{X = x_1\}) + P(\{X = x_2\}) + \dots + P(\{X = x_K\}). \end{aligned}$$

Остаточно одержуємо доведення властивості:

$$p_{k=1} + p_{k=2} + \dots + p_K = 1. \blacktriangle$$

Ряд розподілу дискретної випадкової величини (табл. 4.7), що характеризує всю генеральну сукупність, як правило, через великий обсяг генеральної сукупності не відомий. Одне з завдань вибіркового дослідження в тому й полягає, щоб відповідним чином витягти з генеральної сукупності таку вибірку (показову), вибіркового ряду розподілу відносних частот якої якомога точніше описував би невідомий (генеральний) ряд розподілу.

Розходження між вибіркового рядом розподілу відносних частот і рядом розподілу дискретної випадкової величини можна сформулювати в такий спосіб: вибіркового ряду розподілу відносних частот побудований за вибіркою (за частиною запасу

значень дискретної випадкової величини), а ряд розподілу – за генеральною сукупністю (за повним запасом). Тому вибірковий ряд розподілу відносних частот дає наближений, неповний опис дискретної випадкової величини, а ряд розподілу – її точний, повний опис.

Ряд розподілу дискретної випадкової величини дозволяє обчислювати ймовірності будь-яких подій, пов'язаних із даною дискретною величиною. Як “будівельний матеріал” виступають елементарні події $\{X = x_k\}$, $k = \overline{1, K}$, а “сполучним матеріалом” є дії над подіями (додавання, множення, доповнення, розбиття) і відповідні їм правила (див. розд. 3.1–3.5, 4.1–4.6).

Вибірковий ряд розподілу частот, вибірковий ряд розподілу відносних частот і (генеральний) ряд розподілу дискретної випадкової величини зображуються в такий спосіб. На осі абсцис відкладаються значення груп x_k , від яких вертикально вгору відкладаються відрізки довжини, яка дорівнює частоті, відносній частоті або ймовірності даного значення.

Приклад 4.5 (продовження). Вибірковий ряд розподілу частот для розглядуваної вибірки наведений на рис. 4.6, а. ●

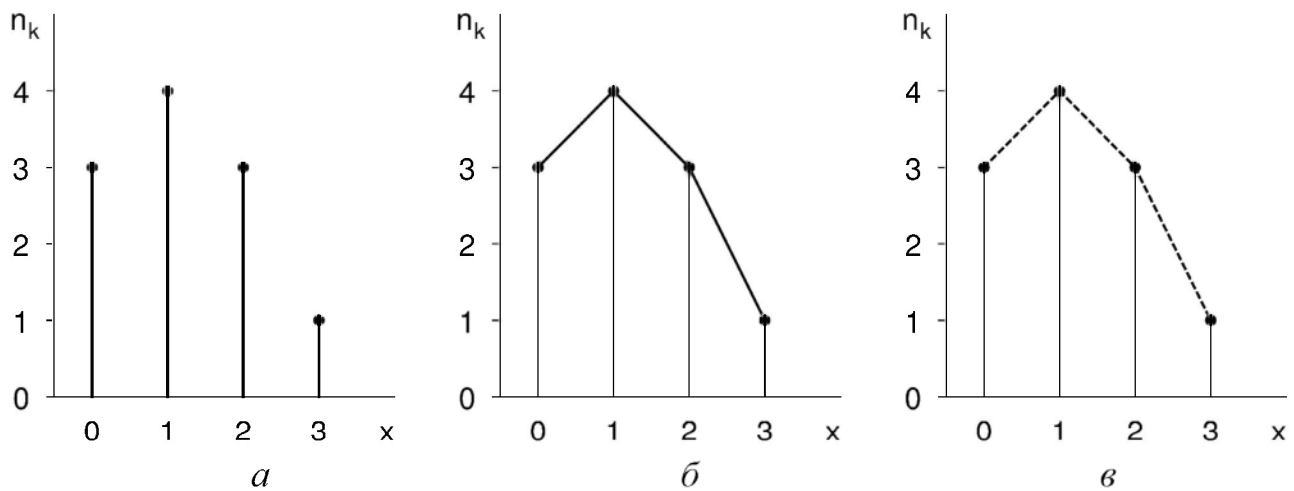


Рис. 4.6. Зображення дискретної випадкової величини: а – вибірковий ряд розподілу частот, б – суцільний полігон частот, в – переривчастий полігон частот

Зауваження 4.8. Іноді при зображенні вибіркового ряду розподілу частот, вибіркового ряду розподілу відносних частот і (генерального) ряду розподілу дискретної випадкової величини кінці вертикальних відрізків з'єднують ламаною. Одержаний “графік” називається **полігоном частот**. Полігон частот створює тільки подобу графіка неперервної функції, оскільки інтервали (x_k, x_{k+1}) є для дискретної випадкової величини неможливими, тому полігон частот краще не використовувати. Якщо ж потрібна наочність у зображенні вибіркового ряду розподілу частот, вибіркового ряду розподілу відносних частот і (генерального) ряду розподілу дискретної випадкової величини, то суцільну ламану в полігоні частот варто замінити переривчастою. ▲

Приклад 4.5 (продовження). Вибіркові ряди розподілу частот у вигляді полігонів частот показані на рис. 4.6, б, в. ●

4.8. Повний опис неперервної випадкової величини

Для неперервної випадкової величини ряд розподілу не може бути складений хоч би тому, що неможливо виписати навіть перший рядок цієї таблиці. Для опису неперервної випадкової величини потрібен інший підхід. Розглянемо деяку вибірку обсягу n : x_i , $i = \overline{1, n}$. Серед вибірових значень знайдемо мінімальне й максимальне:

$$x_{\min} = \min(x_1, x_2, \dots, x_n), \quad x_{\max} = \max(x_1, x_2, \dots, x_n).$$

Виберемо число $a_1 < x_{\min}$ і число $a_{K+1} > x_{\max}$. Причому різниця між $x_{\min} - a_1$ і $a_{K+1} - x_{\max}$ не повинна перевищувати точність вимірів вибірових значень. Тут K – кількість інтервалів, на яку буде розділений відрізок $[a_1, a_{K+1}]$. Кількість інтервалів K не є довільно обраним значенням, а залежить від обсягу вибірки n , тобто $K = K(n)$. Зрозуміло, що немає сенсу обирати його більшим, ніж обсяг вибірки n , але й близьким до n його вибирати не слід, інакше інтервали будуть слабко заповнені вибіровими значеннями і буде неможливо знайти закономірність у їх розподілі.

Кількість інтервалів K , залежно від обсягу вибірки n (рис. 4.7), визначається за однією з таких формул:

$$K_1 = 1 + 3,32 \lg n, \quad ^1 \tag{4.1}$$

$$K_2 = 5 \lg n, \quad ^2 \tag{4.2}$$

$$K_3 = \sqrt{n}. \quad ^3 \tag{4.3}$$

Після вибору кількості інтервалів визначається їх ширина:

$$\Delta a = \frac{a_{K+1} - a_1}{K}.$$

Якщо позначити межі інтервалів через a_k , то за відомих значень a_1 і Δa невідомі межі можна обчислити за рекурентними формулами:

$$a_2 = a_1 + \Delta a, \quad a_3 = a_2 + \Delta a, \quad \dots, \quad a_{k+1} = a_k + \Delta a, \quad \dots, \quad a_{K+1} = a_K + \Delta a.$$

¹ Формула Стерджеса, або Штюржеса (1926).

² Формула Брукса – Краузерса (1963).

³ Формула Доерфеля (1969).

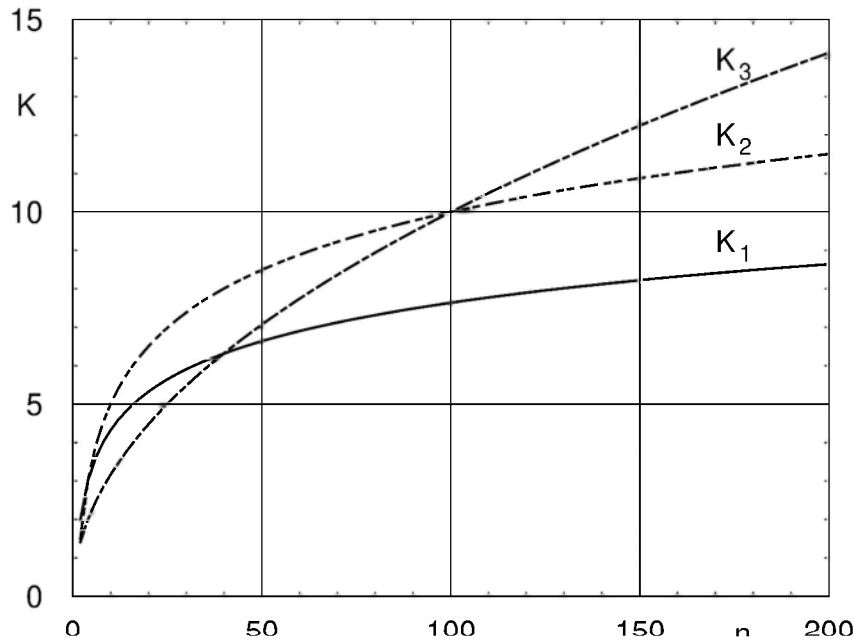


Рис. 4.7. Кількість інтервалів K_1 , K_2 , K_3 залежно від обсягу вибірки n , визначена за формулами (4.1), (4.2), (4.3)

Такий алгоритм уведення інтервалів у загальному випадку не узгоджується з точністю виміру вибірових значень. Справа в тому, що точність, із якою одержані вибірові значення, фактично визначає можливості нашого вимірювального приладу. За допомогою цього вимірювального приладу варто проводити й обробку даних. Нехай, як і в прикладі 4.2, для одержання вихідних даних використовується звичайна лінійка з сантиметровими поділками. Отже, з її допомогою не вдасться відкласти відрізки довжини Δa , заданої в частках сантиметра (десятих, сотих, тисячних і т.д.). Тому результат обчислення Δa округлюється і обчислення меж інтервалів проводиться з округленим значенням Δa . Послідовно відкладаючи відрізки довжини Δa , потрібно стежити за взаємним положенням правої межі останнього побудованого інтервалу і максимального вибірового значення x_{max} . Як тільки чергове значення $a_{k'+1}$ виявиться більшим, ніж x_{max} , побудова інтервалів припиниться. Можливо, що кількість фактично побудованих інтервалів, які включають усі вибірові значення, буде відрізнятись від попередньо знайденого значення K . У такому випадку його варто перевизначити, узявши

$$K + 1 = k' + 1, \text{ тобто } K = k',$$

де k' – номер останнього побудованого інтервалу, що містить максимальне вибірове значення x_{max} . Далі довільний інтервал буде позначатися або своїм номером k , або вказанням на його межі $[a_k, a_{k+1})$.

Зауваження 4.9. Ліва межа включається в інтервал, а права не включається. На це вказують відповідно квадратна та кругла дужки в позначенні інтервалу $[a_k, a_{k+1})$. Причина цієї “нерівноправності” полягає в тому, що довільне вибірове значення x_i

повинне однозначно бути віднесене до одного з інтервалів, а це можна зробити, тільки не включивши в нього одну з меж. Дійсно, якщо вибіркове значення “потрапляє” на межу a_k , то виникає питання: до якого з двох інтервалів – розташованого ліворуч чи праворуч від a_k – його слід віднести. Зазначимо, що зовсім не важливо, яку саме межу – ліву чи праву – включати в інтервал. У нашій роботі буде йти мова про включення лівої межі. ▲

Зауваження 4.10. Зверніть увагу на те, що мала літера k використовується для позначення номера інтервалу, а велика K – для позначення кількості інтервалів. У розд. 4.7 ці літери використовувалися також і для дискретної випадкової величини (див. зауваження 4.4). ▲

Якщо підрахувати кількість вибірових значень для кожного інтервалу, то одержимо **інтервальні частоти** n_k . Їх позначення схожі на ті, які використовуються при роботі з дискретними випадковими величинами, але зрозуміло, що це різні частоти. Множина K пар {інтервали, інтервальні частоти}, тобто $\{[a_k, a_{k+1}), n_k\}$, або {номери інтервалів, інтервальні частоти}, тобто $\{k, n_k\}$, утворює **вибірковий інтервальный ряд розподілу частот** (табл. 4.8).

Таблиця 4.8

Інтервал	$[a_1, a_2)$	$[a_1, a_3)$	...	$[a_k, a_{k+1})$	...	$[a_K, a_{K+1})$
Частота	$n_{k=1}$	$n_{k=2}$	...	n_k	...	$n_{k=K}$

Цей ряд показує, як розподіляються вибірккові значення по інтервалах. Сума інтервальних частот n_k (тобто сума чисел у другому рядку вибіркового інтервального ряду розподілу частот) дорівнює обсягу вибірки n :

$$\sum_{k=1}^n n_k = n.$$

Графічним зображенням вибіркового інтервального ряду розподілу частот є гістограма частот. Це спеціальний графік, що будується на площині в такий спосіб. На осі абсцис відкладаються межі інтервалів a_k . Далі на кожному інтервалі будується прямокутник із висотою, що дорівнює частоті n_k (рис. 4.8).

Припустимо, що обсяг вибірки перевищує значення, за якого виконується умова стабілізації відносних частот для подій

$$\{a_k \leq X < a_{k+1}\}.$$

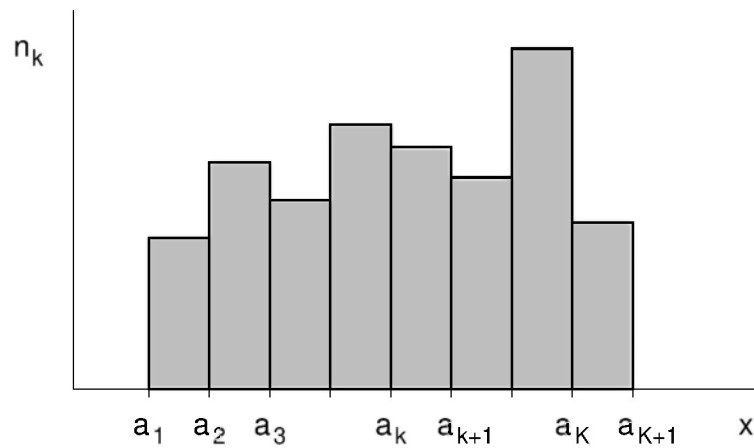


Рис. 4.8. Гістограма частот

У разі збільшення обсягу вибірки n , наприклад, у два рази ($n' = 2n$), “нові” частоти n'_k виявляться приблизно в два рази більшими “старих” частот n_k . Це впливає з властивості стабілізації відносних частот:

$$\frac{n'_k}{n'} \approx \frac{n_k}{n} \Rightarrow n'_k \approx \frac{n'}{n} n_k = 2 n_k.$$

Іншими словами, порівнювати результати двох однотипних серій дослідів різної довжини неможливо.

Вибірковий інтервальний ряд розподілу відносних частот (табл. 4.9) можна одержати з вибіркового ряду розподілу частот (табл. 4.8) діленням частот n_k на обсяг вибірки n .

Таблиця. 4.9

Інтервал	$[a_1, a_2)$	$[a_1, a_3)$	...	$[a_k, a_{k+1})$	...	$[a_K, a_{K+1})$
Відносна частота	$\frac{n_{k=1}}{n}$	$\frac{n_{k=2}}{n}$	...	$\frac{n_k}{n}$	...	$\frac{n_{k=K}}{n}$

Сума відносних частот у другому рядку вибіркового інтервального ряду розподілу відносних частот дорівнює 1:

$$\sum_{k=1}^K \frac{n_k}{n} = \frac{n_1}{n} + \frac{n_2}{n} + \dots + \frac{n_K}{n} = \frac{n_1 + n_2 + \dots + n_K}{n} = \frac{n}{n} = 1.$$

Гістограма відносних частот будується аналогічно гістограмі частот. Цілком очевидно, що відносні частоти інтервалів не повинні “відгукуватися” на зміну обсягу вибірки, проте за різних обсягів вибірок кількість інтервалів і їх ширина розріз-

няються. Виникає питання, яким у такому випадку буде вибірковий інтервальний ряд розподілу відносних частот.

Розглянемо відгук вибіркового інтервального ряду розподілу відносних частот тільки на зміну ширини інтервалу у разі штучного обмеження обсягу вибірки. Збільшення ширини інтервалу у два рази можна змодельювати об'єднанням двох суміжних інтервалів (рис. 4.9). При об'єднанні інтервалів шириною Δa з частотами n_k і n_{k+1} утворюється інтервал шириною $2\Delta a$ з частотою $(n_k + n_{k+1})$ і відносною частотою, що дорівнює сумі відносних частот об'єднаних інтервалів:

$$\frac{n_k + n_{k+1}}{n} = \frac{n_k}{n} + \frac{n_{k+1}}{n}.$$

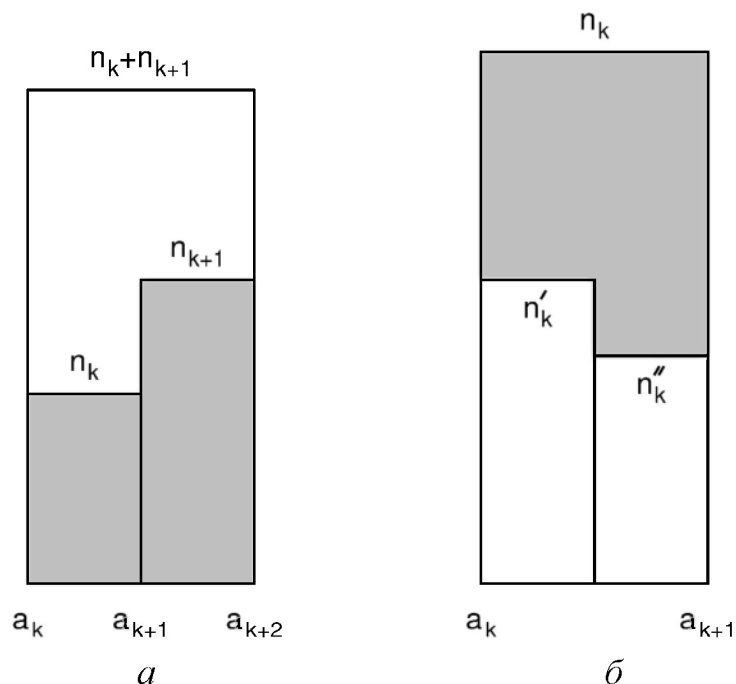


Рис. 4.9. Відгук гістограми частот: a – на об'єднання суміжних інтервалів; b – на розділення інтервалу (висота “нового” стовпчика (білий колір) при об'єднанні інтервалів дорівнює сумі висот “старих” (сірий колір), висоти “нових” стовпчиків (білий колір) при розділенні інтервалу в сумі дорівнюють висоті “старого” (сірий колір))

Для поділу інтервалу ширини Δa на два інтервали ширини $\Delta a/2$ візьмемо, що частоти “нових” інтервалів

$$n'_k \approx n''_k \approx \frac{1}{2} n_k,$$

тоді їх відносні частоти зменшаться приблизно у два рази:

$$\frac{n'_k}{n} \approx \frac{n''_k}{n} \approx \frac{1}{2} \frac{n_k}{n}.$$

Отже, порівнювати дві серії дослідів, оброблених за допомогою вибіркового інтервального ряду розподілу відносних частот, що мають різну ширину інтервалів, також неможливо.

Уведемо на інтервалах $[a_k, a_{k+1})$ замість відносних частот нову величину, названу **вибірковою функцією густини**, значення якої на інтервалі k дорівнює

$$f_k^* = \frac{p_k^*}{\Delta a}. \quad (4.4)$$

Множина K пар $\{k, f_k^*\}$ утворює **вибірковий інтервальный ряд розподілу функції густини** (табл. 4.10).

Таблиця 4.10

Інтервал	$[a_1, a_2)$	$[a_1, a_3)$	...	$[a_k, a_{k+1})$	...	$[a_K, a_{K+1})$
Вибіркова функція густини	$f_{k=1}^*$	$f_{k=2}^*$	...	f_k^*	...	$f_{k=K}^*$

Сума значень вибіркової функції густини в другому рядку табл. 4.10 дорівнює оберненій величині ширини інтервалу:

$$\sum_{k=1}^n f_k^* = \sum_{k=1}^n \frac{1}{\Delta a} \frac{n_k}{n} = \frac{1}{\Delta a} \sum_{k=1}^n \frac{n_k}{n} = \frac{1}{\Delta a} \frac{1}{n} \sum_{k=1}^n n_k = \frac{1}{\Delta a} \frac{n}{n} = \frac{1}{\Delta a}.$$

Гістограма вибіркової функції густини будується аналогічно гістограмам частот і відносних частот.

Легко помітити, що за сталої ширини інтервалів вибіркового інтервального ряду розподілу функції густини “не реагує” на зміну обсягу вибірки, тому що в чисельнику виразу для f_k^* знаходиться відносна частота, для якої виконується умова стабілізації. Після об'єднання двох суміжних інтервалів значення вибіркової функції густини “об'єднаного” інтервалу дорівнює

$$f_{k,k+1}^* = \frac{\frac{n_k + n_{k+1}}{n}}{2\Delta a} = \frac{1}{2} \frac{\frac{n_k}{n} + \frac{n_{k+1}}{n}}{\Delta a} = \frac{1}{2} \left(\frac{p_k^*}{\Delta a} + \frac{p_{k+1}^*}{\Delta a} \right) = \frac{1}{2} (f_k^* + f_{k+1}^*),$$

тобто середньому арифметичному значенню вибірових функцій густини об'єднаних інтервалів, отже, воно є обмежене значеннями f_k^* і f_{k+1}^* (рис. 4.10).

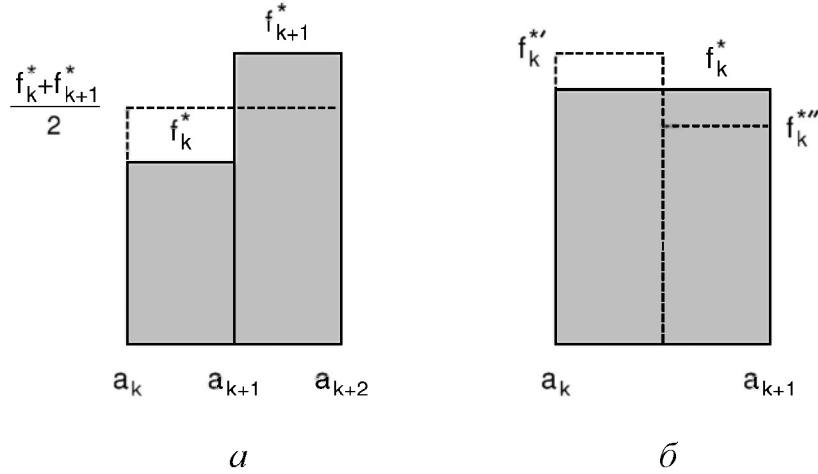


Рис. 4.10. Відгук гістограми вибіркової функції густини: *a* – на об’єднання суміжних інтервалів; *б* – на розділення інтервалу (стовпчики до об’єднання і до розділення показані сірим кольором, межі стовпчиків після об’єднання й після розділення проведені пунктиром)

Для поділу інтервалу шириною Δa на два інтервали шириною $\Delta a/2$ візьмемо, що частоти й відносні частоти нових інтервалів зв’язані співвідношеннями

$$n'_k \approx n''_k \approx \frac{1}{2} n_k, \quad \frac{n'_k}{n} \approx \frac{n''_k}{n} \approx \frac{1}{2} \frac{n_k}{n},$$

звідки впливають такі оцінки для значень вибіркової функції густини “нових” інтервалів:

$$f_k^{*'} = \frac{n'_k / n}{\Delta a / 2} = \frac{2 n'_k / n}{\Delta a} \approx \frac{n_k / n}{\Delta a} = f_k^*,$$

$$f_k^{*''} = \frac{n''_k / n}{\Delta a / 2} = \frac{2 n''_k / n}{\Delta a} \approx \frac{n_k / n}{\Delta a} = f_k^*,$$

тобто значення вибіркової функції густини на “нових” інтервалах приблизно дорівнюють значенням вибіркової функції густини “старого” інтервалу.

На вибіркового інтервального ряді розподілу функції густини не позначається ні зміна обсягу вибірки, ні зміна ширини інтервалу. Тобто у вигляді вибіркової функції густини вдалося одержати **інваріантну** величину для обробки результатів різних серій однотипних дослідів.

Графік вибіркової функції густини $f^*(x)$ можна одержати, якщо на гістограмі вибіркової функції густини f_k^* забрати стовпчики, а залишити тільки верхні їх сторони (рис. 4.11).

Граничне положення вибіркової функції густини $f^*(x)$, визначеної на інтервалах $k = \overline{1, K}$, за необмеженого збільшення обсягу вибірки n і зменшення до нуля ширини інтервалів (за необмеженого збільшення кількості інтервалів) називається **генеральною функцією густини неперервної випадкової величини**. Для функції густини неперервної випадкової величини X використовується позначення $f(x)$. Графік функції густини $f(x)$ наведений на рис. 4.11.

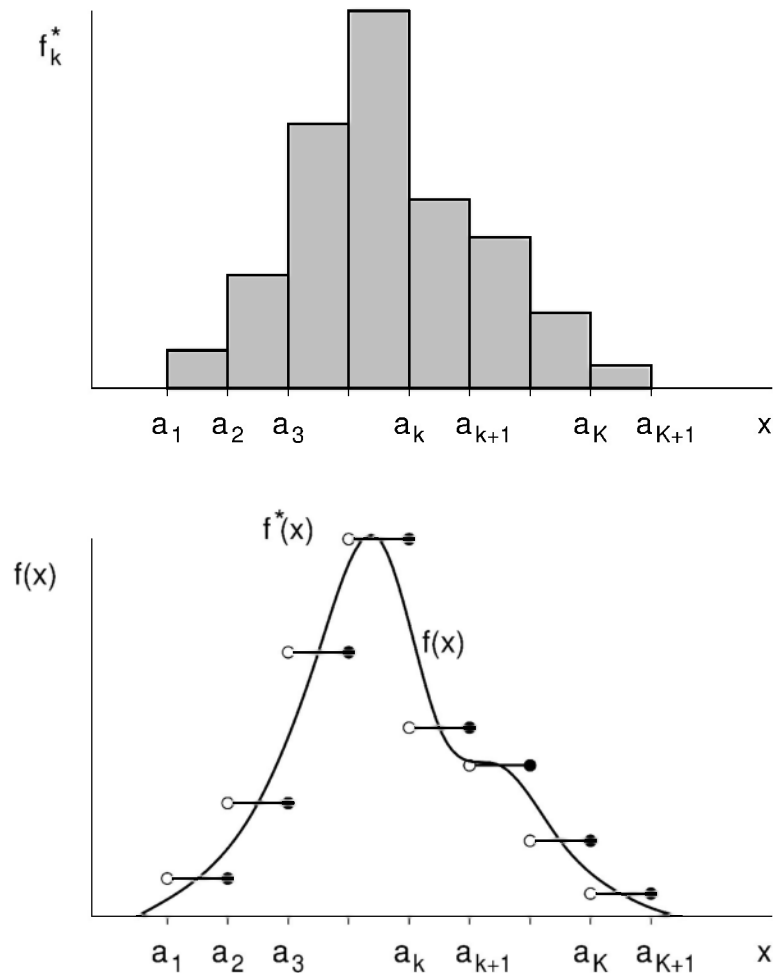


Рис. 4.11. Гістограма вибіркової функції густини f_k^* (вгорі), графік вибіркової функції густини $f^*(x)$ (верхні сторони стовпчиків гістограми, унизу), графік функції густини $f(x)$ (суцільна лінія, унизу)

Зауваження 4.11. Область визначення функції густини $f(x)$ завжди обмежена. Це впливає з обмеженості результатів вимірів при вивченні будь-якого реального процесу (не обов'язково випадкового). Але з досліду ми знаємо тільки мінімальне й максимальне вибіркові значення, а мінімальне і максимальне генеральні значення, тобто межі області визначення $f(x)$, як правило, невідомі. Зрозуміло, що працювати з функцією, область визначення якої не відома, не завжди зручно, а іноді – неможливо. Тому додатково визначають функцію $f(x)$ у нескінченній області ліворуч від мінімального значення (тобто від $-\infty$ до мінімального значення) і в нескінченній об-

ласті праворуч від максимального значення (тобто від максимального значення до $+\infty$) нульовими значеннями. У результаті функція густини $f(x)$ виявляється визначеною на всій числовій осі $(-\infty, +\infty)$. Те, що в область визначення включені фізично неможливі значення, на одержуваних результатах ніяк не позначається, адже рівність нулеві функції густини в будь-якому інтервалі означає, що значень із цього інтервалу випадкова величина X ніколи не набуває (іншими словами, ця подія неможлива). ▲

Властивості функції густини

1. $f(x) \geq 0$, тобто функція густини невід'ємна.

Розглянемо вибіркового інтервального ряд розподілу функції густини. За визначенням (4.4) на кожному інтервалі k значення вибіркової функції густини дорівнюють

$$f_k^* = \frac{p_k}{\Delta a}.$$

Оскільки ширина інтервалу відмінна від нуля, тобто $\Delta a > 0$ (знаменник), а ймовірність набуття випадковою величиною значень у довільному інтервалі невід'ємна, тобто $p_k^* \geq 0$ (чисельник), то $f_k^* \geq 0$ незалежно від значень n і Δa . Отже, при граничному переході від f_k^* до $f(x)$ властивість невід'ємності зберігається.

2. Площа, обмежена віссю абсцис і графіком функції густини $f(x)$ (для простоти будемо говорити про площу під кривою $f(x)$) дорівнює 1.

Площа, яку займає гістограма вибіркової функції густини, дорівнює сумі площ складових її частин (прямокутників) (див. рис. 4.11). Площа S_k прямокутника, що належить до інтервалу k , дорівнює

$$S_k = \Delta a \cdot f_k^*,$$

де Δa – ширина прямокутника, а f_k^* – його висота.

Отже, площа гістограми дорівнює

$$\begin{aligned} S &= S_1 + S_2 + \dots + S_K = \\ &= \Delta a \cdot f_1^* + \Delta a \cdot f_2^* + \dots + \Delta a \cdot f_K^* = \Delta a (f_1^* + f_2^* + \dots + f_K^*). \end{aligned}$$

Сума значень вибіркової функції густини для всіх інтервалів, що входить у вираз для S , дорівнює $1/\Delta a$. Тому

$$S = \Delta a (f_1^* + f_2^* + \dots + f_K^*) = \frac{\Delta a}{\Delta a} = 1.$$

Як бачимо, результат не залежить від значень n і Δa . Тому він може бути поширений і на функцію густини. Іноді його записують у такій формі:

$$\int_{-\infty}^{+\infty} f(x) dx = 1,$$

що ґрунтується на геометричному тлумаченні визначеного інтеграла й розширенні області визначення $f(x)$ (див. зауваження 4.11). ■

3. Ймовірність набуття випадковою величиною X значень у деякому інтервалі $[a, b)$ дорівнює площі S під кривою функції густини $f(x)$, що відтинається вертикальними прямими $x = a$ і $x = b$ (рис. 4.12).

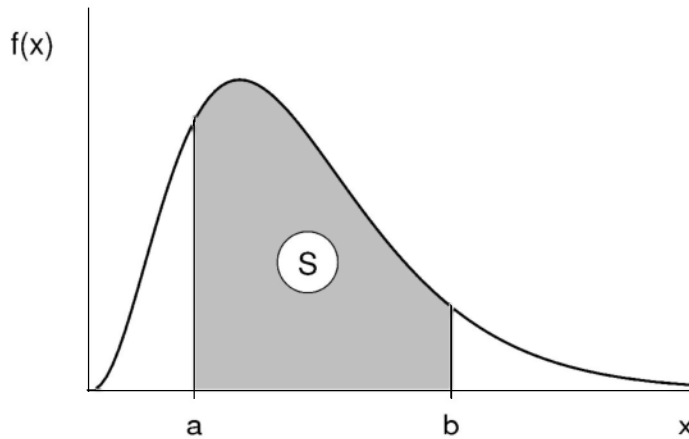


Рис. 4.12. Схема до формулювання властивості 3 функції густини $f(x)$

Для обґрунтування властивості розглянемо кілька окремих задач.

Як відомо, ймовірність набуття випадковою величиною значень у деякому інтервалі k за умови стабілізації відносних частот дорівнює p_k^* . Цю саму ймовірність можна виразити через значення вибіркової функції густини на даному інтервалі або через площу відповідного прямокутника гістограми $p_k^* = \Delta a \cdot f_k^* = S_k$ (див. властивість 2). Ймовірність набуття випадковою величиною значень у декількох послідовно розташованих інтервалах буде дорівнювати сумі площ прямокутників, що відповідають даним інтервалам.

Якщо за побудованою для вибірки, яка має обсяг n' і кількість інтервалів K' , гістограмою вибіркової функції густини підібрати положення меж інтервалів a' і b' , найбільш близьких до заданих значень a і b , то шукану ймовірність $P(\{a \leq X < b\})$ можна визначити наближено як сумарну площу S' прямокутників гістограми, розташованих між підібраними межами інтервалів a' і b' . Поліпшити оцінку для $P(\{a \leq X < b\})$ можна шляхом збільшення обсягу вибірки до $n'' > n'$ і, відповідно, кількості інтервалів $K'' > K'$. За допомогою нової гістограми, побудованої за вибір-

кою більшого обсягу, можна підібрати межі інтервалів a'' і b'' , розташованих ближче до заданих значень a і b (рис. 4.13). У цьому випадку розбіжність між величинами площ S і S'' буде меншою, ніж між величинами площ S і S' :

$$|S - S''| < |S - S'|.$$

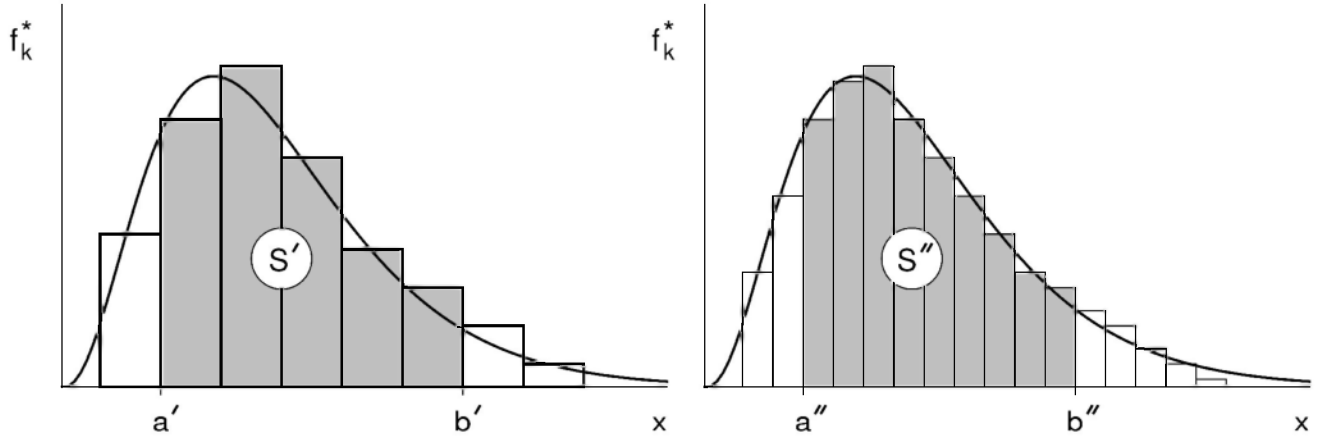


Рис. 4.13. Схема до обґрунтування властивості 3 функції густини $f(x)$

Продовжуючи підбір до вичерпання вибіркою генеральної сукупності, дійдемо висновку, що шукана ймовірність буде дорівнювати сумарній площі групи прямокутників гістограми, відібраних за принципом близькості лівої межі крайнього лівого прямокутника до значення a і близькості правої межі крайнього правого прямокутника до значення b (рис. 4.13). У даному випадку сумарна площа цих дуже вузьких прямокутників (їх ширина Δa наближається до нуля) буде дорівнювати площі під кривою функції густини $f(x)$ (верхні частини прямокутників, що утворюють графік $f^*(x)$, гранично переходять у графік $f(x)$, тобто $f^*(x) \rightarrow f(x)$), а це означає, що властивість доведена. Математично вона записується так:

$$S = P(\{a \leq X < b\}) = \int_a^b f(x) dx. \blacksquare$$

4. Ймовірність того, що випадкова величина X набуває значень, менших заданого значення x , дорівнює площі S_1 під кривою функції густини ліворуч від значення x (рис. 4.14).

Ця властивість виводиться з властивості 3 за умови $a = -\infty$, $b = x$.

5. Ймовірність того, що випадкова величина X набуває значень, більших ніж задане значення x , дорівнює площі S_2 під кривою функції густини праворуч від значення x (рис. 4.14).

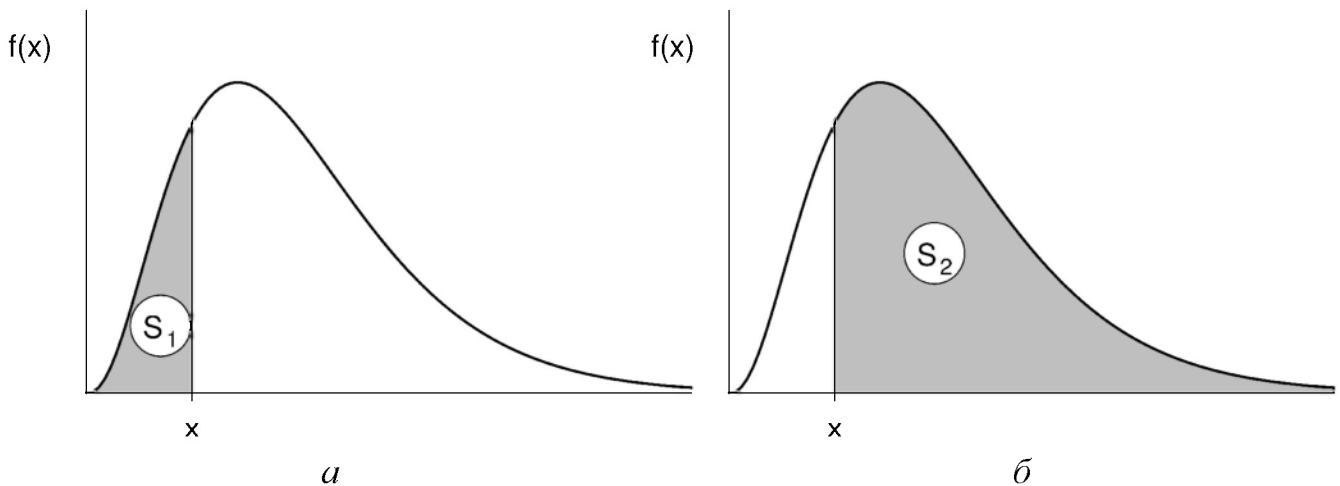


Рис. 4.14. Схема до обґрунтування властивостей 4 (а) і 5 (б) функції густини $f(x)$

Ця властивість виводиться з властивості 3 за умови $a = x$, $b = +\infty$. Можна також застосувати висновок із правила додавання для двох несумісних подій.

Наведені властивості функції густини не є вичерпні, але формулювання інших подамо у розд. 4.11.

4.9. Порівняння понять функції густини та фізичної густини

З'ясуємо, що мають на увазі, коли говорять про функцію густини: чи це та сама величина, яка вводиться в курсі фізики?

У фізиці густина вводиться в такий спосіб. Нехай є деяке просторове тіло, наприклад цеглина, з відомими масою M і об'ємом V (рис. 4.15).

За визначенням **середньої густини** матеріалу (речовини) даного тіла називається відношення

$$\rho = \frac{M}{V}.$$

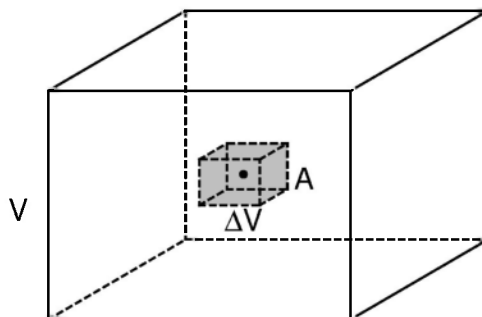


Рис. 4.15. Схема до визначення об'ємної густини маси

Якщо матеріал тіла не є однорідний, то ця характеристика мало про що буде говорити. Бажано знати, наскільки неоднорідним є матеріал. Для цього вводиться по-

няття локальної густини. Нехай якийсь окіл деякої точки A тіла (цей окіл – маленький шматочок тіла досить довільної форми, для простоти – маленький паралелепіпед, що містить усередині дану точку) має масу ΔM й об'єм ΔV (рис. 4.15). Зрозуміло, що за аналогією з середньою густиною матеріалу всього тіла можна ввести й середню густину матеріалу в околі вибраної точки A . Якщо зменшувати розмір цього околу (шматочка тіла), то можна перейти до поняття густини матеріалу в даній точці. Отже, **об'ємна густина маси** в даній точці дорівнює

$$\rho_A = \lim_{\Delta V \rightarrow 0} \frac{\Delta M}{\Delta V}.$$

У цьому визначенні мова не йде про границю в тому розумінні, яке прийняте в курсі вищої математики. Там розглядається границя цілком визначеної послідовності або цілком визначеної функції деякої змінної. Тут же відсутні визначені в цьому сенсі як послідовність, так і функція. Мається на увазі таке: якщо довільно зменшувати поперечні розміри околу точки A (маленького шматочка матеріалу), то, починаючи з деякого характерного розміру, значення відношення ΔM до ΔV з прийнятою точністю перестане змінюватися. Однак зауважимо, що за довільного зменшення цілком визначена числова послідовність $\frac{\Delta M}{\Delta V}$ відсутня.

Із визначення об'ємної густини маси випливає, що якщо в даній точці A відома величина ρ_A , то масу довільного за формою маленького шматочка тіла з об'ємом ΔV , який містить усередині точку A , наближено можна визначити в такий спосіб:

$$\Delta M \approx \rho_A \Delta V,$$

причому наближена рівність виконується тим точніше, чим менший об'єм шматочка ΔV .

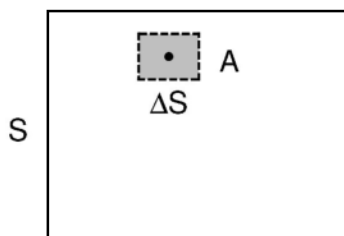


Рис. 4.16. Схема до визначення поверхневої густини маси

Розглянемо тіло, одним із розмірів якого можна знехтувати порівняно з двома іншими (наприклад, аркуш паперу, лист фанери, дахового заліза і т.ін. (рис. 4.16)).

Якщо маса листа такого матеріалу дорівнює M , а його площа – S , то середньою поверхневою густиною матеріалу варто вважати відношення

$$\rho = \frac{M}{S}.$$

Для неоднорідного матеріалу можна ввести локальну характеристику розподілу матеріалу в кожній точці, якщо розглянути окіл точки з масою ΔM і площею ΔS . Форма цього околу не має значення, але можна його зобразити у вигляді маленького прямокутника, що містить усередині дану точку A . Тоді густина матеріалу в даній точці

$$\rho_A = \lim_{\Delta S \rightarrow 0} \frac{\Delta M}{\Delta S}.$$

Уведену в такий спосіб густину називають **поверхневою густиною маси**. Це поняття можна застосовувати приблизно так само, як і поняття об'ємної густини маси. Якщо в даній точці A відома величина ρ_A , то масу довільного за формою маленького шматочка листового матеріалу з площею ΔS , що містить усередині точку A , наближено можна визначити в такий спосіб:

$$\Delta M \approx \rho_A \Delta S,$$

причому наближена рівність виконується тим точніше, чим менша площа шматочка ΔS .

Нехай є тіло, двома розмірами якого можна знехтувати порівняно з третім, наприклад шматок труби, проводу і т.ін. (рис. 4.17). Якщо маса шматка такого матеріалу дорівнює M , а його довжина – L , то середньою лінійною густиною матеріалу варто вважати відношення

$$\rho = \frac{M}{L}.$$

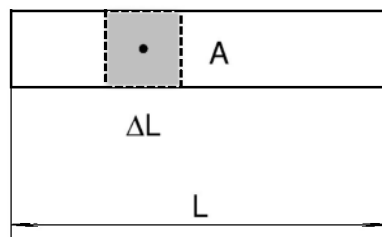


Рис. 4.17. Схема до визначення лінійної густини маси

Для неоднорідного матеріалу можна ввести локальну характеристику розподілу матеріалу в кожній точці, якщо розглянути окіл даної точки A з масою ΔM і довжиною ΔL . Цей окіл являє собою маленький відрізок розглянутого шматка труби, що містить усередині дану точку A . Тоді густина матеріалу в даній точці

$$\rho_A = \lim_{\Delta L \rightarrow 0} \frac{\Delta M}{\Delta L}.$$

Уведену в такий спосіб густину називають **лінійною густиною маси**. Це поняття можна застосовувати приблизно так само, як і поняття просторової та поверхневої густини маси. Тобто якщо в даній точці A відома величина ρ_A , то масу довільного маленького шматочка матеріалу довжиною ΔL , що містить усередині точку A , наближено можна визначити в такий спосіб:

$$\Delta M \approx \rho_A \Delta L.$$

З'ясуємо, що спільного в цих трьох визначеннях густини. Зверніть увагу на те, що це визначення **густини маси**. Саме відношення маси маленького шматочка матеріалу до міри цього шматочка ("міра" – це спеціальний математичний термін, тут він застосовується для узагальненого позначення об'єму, площі або довжини шматочка) і дає один із видів густини маси.

Проаналізуємо, як саме стосується поняття густини як такої до функції густини. Звернемо увагу на те, що при введенні вибіркової функції густини для довільного інтервалу (статистична) ймовірність того, що значення досліджуваної величини знаходяться в межах цього інтервалу, поділяється на його міру, тобто довжину. У даному випадку мова йде про **густину ймовірності**, а точніше про **лінійну густину ймовірності**. При переході від вибіркової функції густини $f^*(x)$ до звичайної функції густини $f(x)$ робиться граничний перехід, подібний до переходу від середньої лінійної густини маси до лінійної густини в точці. Але ми вилучили слово "ймовірність" із визначення вибіркової функції густини $f^*(x)$ і функції густини $f(x)$.

В науковій літературі нечасто зустрічається вираз "густина маси" – зрозуміло, що мова йде про масу, тому це слово опускають. Однак і в математиці, і в фізиці зручно вводити поняття густини для різних величин, отже, в кожному конкретному випадку варто вказувати, яка саме величина мається на увазі. Ми будемо вести мову виключно про густину ймовірності, тому слово "ймовірність", як правило, буде опускатися. Однак не завжди густина ймовірності буде лінійною.

Поняття функції густини (ймовірності) $f(x)$ застосовується в такий спосіб. Нехай у деякій точці $x = c$ відоме значення $f(c)$. Тоді ймовірність того, що випадкова величина X набуває значень у деякому околі цієї точки, тобто в малому інтервалі (a, b) довжиною $\Delta x = b - a$, що містить точку c , або, іншими словами, того, що в експерименті будуть зареєстровані значення в цьому інтервалі, приблизно дорівнює

$$P(\{a \leq X < b\}) \approx f(c) \cdot (b - a) = f(c) \Delta x, \quad c \in (a, b).$$

Дана ймовірність дорівнює площі S фігури під графіком $f(x)$, що відтинається двома вертикальними прямими $x=a$ і $x=b$ (рис. 4.18). Точне значення цієї площі, відповідно до властивості 3, дорівнює

$$S = \int_a^b f(x) dx,$$

і за теоремою про середнє значення може бути подане як

$$S = f(\xi) \cdot (b - a), \quad \xi \in (a, b).$$

Якщо взяти $\xi \approx c$, вийде наближена рівність, записана вище.

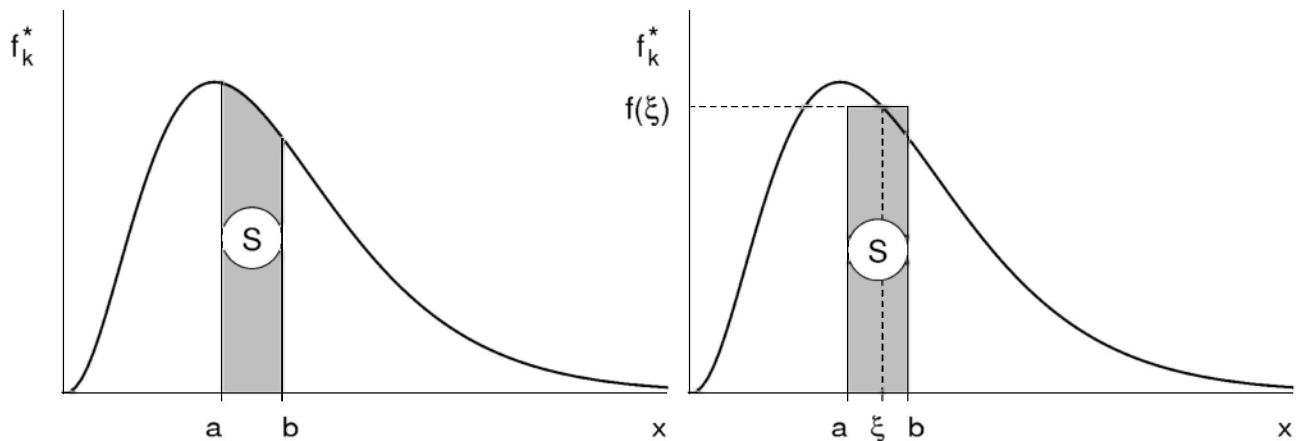


Рис. 4.18. Схема до тлумачення функції густини $f(x)$

4.10. Функція розподілу випадкової величини

Функція розподілу застосовується для повного опису випадкової величини. Поняття ряду розподілу дискретної випадкової величини й функції густини неперервної випадкової величини є спеціальними способами повного опису, застосовними тільки для одного типу випадкової величини. Функція розподілу вводить більш загальний спосіб повного опису, застосовний як для дискретної, так і для неперервної випадкової величини.

Функція розподілу $F(x)$ випадкової величини X дорівнює ймовірності події, яка полягає в тому, що випадкова величина X набуває значення, меншого, ніж задане x :

$$F(x) = P(\{X < x\}).$$

Подія, про яку йде мова у визначенні, наведена на рис. 4.19.

Зверніть увагу на те, як використовується літера (X, x) у визначенні функції розподілу: мала – для значень випадкової величини, у тому числі як аргумент (залежна змінна) функції розподілу, а велика – для самої випадкової величини (див. розд. 4.2).

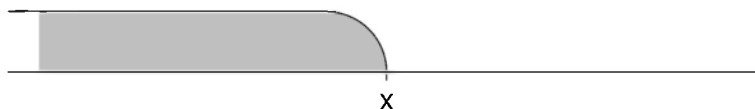


Рис. 4.19. Схема до визначення функції розподілу $F(x)$

Властивості функції розподілу

1. Значення $F(x)$ знаходяться в межах 0 і 1:

$$0 \leq F(x) \leq 1.$$

Оскільки значення функції розподілу дорівнює ймовірності відповідної події, то, враховуючи визначення ймовірності, можна вважати, що властивість доведена. ■

2. Функція $F(x)$ є невід'ємною функцією свого аргументу. Це означає, що для будь-яких двох значень a, b випадкової величини, таких, що $b > a$, виконується нерівність

$$F(b) \geq F(a).$$

Для обґрунтування властивості введемо три допоміжні події (рис. 4.20):

$$\begin{aligned} A &= \{X < a\}; \\ B &= \{a \leq X < b\}; \\ C &= \{b \leq X\}, \end{aligned}$$

причому $AB = \emptyset$, $AC = A$, $BC = B$, $C = A + B$

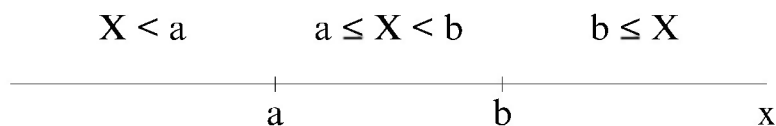


Рис. 4.20. Схема до обґрунтування властивості 2 функції розподілу $F(x)$

За правилом додавання для несумісних подій (див. розд. 4.4)

$$P(C) = P(A + B) = P(A) + P(B),$$

або

$$P(\{X < b\}) = P(\{X < a\}) + P(\{a \leq X < b\}).$$

Якщо застосувати визначення функції розподілу, то правило додавання можна записати так:

$$F(b) = F(a) + P(\{a \leq X < b\})$$

або так:

$$F(b) - F(a) = P(\{a \leq X < b\}).$$

Оскільки відповідно до властивості 1 функції розподілу $P(\{a \leq X < b\}) \geq 0$, то

$$F(b) - F(a) \geq 0 \Rightarrow F(b) \geq F(a). \blacksquare$$

3. $F(b) - F(a) = P(\{a \leq X < b\})$. Ця властивість є проміжним результатом при доведенні властивості 2.

4. $F(-\infty) = 0$.

Оскільки з визначення функції розподілу випливає, що:

$$F(x_{\min}) = 0,$$

то, з урахуванням зауваження 4.11, властивість доведена. $\blacksquare$

5. $F(+\infty) = 1$.

Оскільки праворуч від $x_{\max}$ немає значень випадкової величини, то виконується нерівність $P(\{X > x_{\max}\}) = 0$. $\blacksquare$

Зв'язок функції розподілу та функції густини неперервної випадкової величини

Із властивостей 3 функції розподілу та функції густини (див. розд. 4.8) випливає ключове рівняння, що зв'язує функцію розподілу й функцію густини для неперервної випадкової величини:

$$\int_a^b f(x) dx = F(b) - F(a).$$

Це формула Ньютона – Лейбніца з курсу вищої математики. Якщо функція розподілу $F(x)$ диференційована, то з формули Ньютона – Лейбніца випливає простий вираз зв'язку функції розподілу та функції густини:

$$f(x) = F'(x).$$

Ці два вирази можна вважати додатковими властивостями функції розподілу й функції густини.

Вибірковий спосіб уведення функції розподілу для неперервної випадкової величини

Уведемо поняття функції розподілу для неперервної випадкової величини в спосіб, відмінний від викладеного в попередньому розділі і подібний до способу введення поняття функції густини.

Розглянемо вибіркові дані x_i , $i = \overline{1, n}$, розподілені за K інтервалами у вигляді вибіркового інтервального ряду розподілу частот, вибіркового інтервального ряду розподілу відносних частот і вибіркового інтервального ряду розподілу функції густини (див. розд. 4.8). Уведемо вибіркову функцію $F^*(x)$, спочатку визначивши її на межах інтервалів таким чином:

$$F^*(a_k) = P^*(\{X < a_k\}), \quad k = \overline{1, K}.$$

Розпишемо це визначення для меж інтервалів:

$$F^*(a_1) = P^*(\{X < a_1\}) = 0;$$

$$F^*(a_2) = P^*(\{X < a_2\}) = P^*(\{X < a_1\}) + P^*(\{a_1 \leq X < a_2\}) = p_{k=1}^*;$$

$$\begin{aligned} F^*(a_3) &= P^*(\{X < a_3\}) = P^*(\{X < a_2\}) + P^*(\{a_2 \leq X < a_3\}) = \\ &= p_{k=1}^* + p_{k=2}^* \end{aligned}$$

.....

$$\begin{aligned} F^*(a_K) &= P^*(\{X < a_K\}) = P^*(\{X < a_{K-1}\}) + P^*(\{a_{K-1} \leq X < a_K\}) = \\ &= p_{k=1}^* + p_{k=2}^* + \dots + p_{k=K-1}^* \end{aligned}$$

$$\begin{aligned} F^*(a_{K+1}) &= P^*(\{X < a_{K+1}\}) = P^*(\{X < a_K\}) + P^*(\{a_K \leq X < a_{K+1}\}) = \\ &= p_{k=1}^* + p_{k=2}^* + \dots + p_{k=K-1}^* + p_{k=K}^* = 1. \end{aligned}$$

Додатково визначимо функцію в середині інтервалів. Розв'язок цієї задачі не єдиний, оскільки групування вибірових значень за інтервалами позбавляє індивідуальності всі значення, що потрапляють у певний інтервал (див. розд. 4.3, 4.8). Тому будемо вважати функцію $F^*(x)$ сталою в середині інтервалів $[a_k, a_{k+1})$ (див. рис.4.19), причому

$$F_k^* = F^*(a_{k+1}), \quad k = \overline{1, K}.$$

Одержану функцію будемо називати **вибірковою функцією розподілу** $F^*(x)$ неперервної випадкової величини X . Її графік (рис.4.19), складається з горизонтальних відрізків, що належать до відповідних інтервалів, причому лівий кінець відрізка не належить до даного інтервалу (показаний білим кружечком), а правий – належить (показаний чорним кружечком).

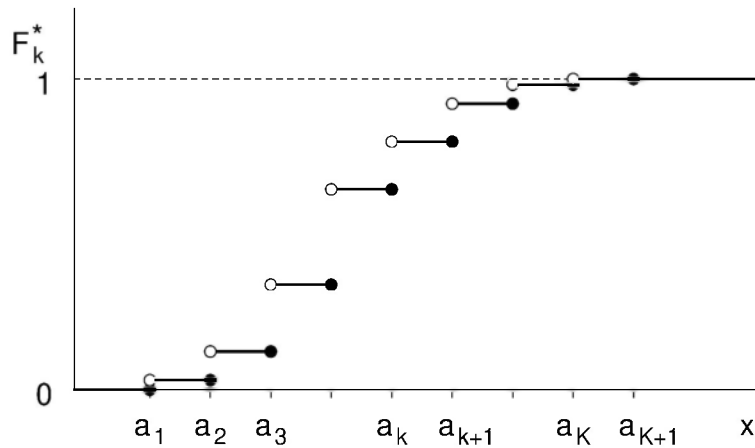


Рис. 4.21. Графік вибіркової функції розподілу

Легко помітити, що:

$$F^*(x) \equiv 0, \quad x < a_1,$$

$$F^*(x) \equiv 1, \quad x > a_{K+1}.$$

Границя вибіркової функції розподілу за виконання умов $\Delta a \rightarrow 0$ ($K \rightarrow \infty$) і $n \rightarrow N(\infty)$ називається **функцією розподілу** неперервної випадкової величини X (рис. 4.21).

4.11. Вибіркові числові характеристики

Крім повного (наближеного або точного) опису випадкової величини у вигляді вибіркового ряду розподілу відносних частот і (генерального) ряду розподілу дискретної випадкової величини та вибіркового інтервального ряду розподілу відносних частот або (генеральної) функції густини неперервної випадкової величини, існує ще й спосіб неповного опису. Його сутність полягає в тому, що розподіл випадкової величини (ймовірність набуття дискретною випадковою величиною певних значень або неперервною випадковою величиною значень у певних інтервалах) замінюють декількома числами. Це подібне до середньої густини шматка матеріалу (див. розд. 4.10), діапазону зміни температури атмосферного повітря за рік і т.д.

Уведемо кілька таких показників, що задають неповний опис, припускаючи, що існує вибірка x_i , $i = \overline{1, n}$.

Вибірковим середнім значенням випадкової величини X називається така величина:

$$m = m_X = \frac{x_1 + x_2 + \dots + x_n}{n} = \frac{1}{n} \sum_{i=1}^n x_i. \quad (4.5)$$

Це є **арифметичне середнє** чисел x_i , $i = \overline{1, n}$.

Зауваження 4.12. Якщо в досліді вивчається тільки одна випадкова величина (ознака), то можна рівноправно використовувати позначення m і m_X . У дослідях із декількома випадковими величинами X , Y і т. д. для розрізнення вибірових середніх значень використовують позначення з нижнім індексом: m_X , m_Y і т. ін. ▲

Зауваження 4.13. Зверніть увагу на те, що розмірність вибірового середнього значення випадкової величини така сама, що й самої випадкової величини, тобто $[m_X] = [X]$. Наприклад, якщо значення випадкової величини вимірюються в кілограмах, метрах, відсотках і т.д., то до обчисленого за результатами серії n дослідів вибірового середнього значення (4.5) варто приписати праворуч таку ж розмірність. ▲

Зазначимо, що вибірове середнє значення може дати занадто обмежену інформацію про розподіл вибірових значень.

Приклад 4.6. Дві дискретні випадкові величини X і Y задані своїми вибіровими рядами розподілу частот, наведеними в табл. 4.11 і 4.12.

Таблиця 4.11

x_k	-2	1	5	9	12
n_k	1	10	25	10	1

Таблиця 4.12

y_l	0	4	8
n_l	20	10	20

Легко переконатися, що дискретні випадкові величини X і Y мають різні множини значень, проте їх вибірові середні значення збігаються: $m_X = m_Y = 5$. ●

Зауваження 4.3. (продовження). Якщо в досліді вивчається декілька випадкових величин, то пари літер (i, n) буде недостатньо для позначення номера вибірового значення й обсягу вибірки, а однієї літери (k, K) – недостатньо для позначень при записі вибірового ряду розподілу частот, вибірового ряду розподілу відносних частот або (генерального) ряду розподілу дискретної випадкової величини і вибірового інтервального ряду розподілу частот, вибірового інтервального ряду розподілу відносних частот або (генеральної) функції густини неперервної випадкової величини. У таких дослідях будемо використовувати додаткові літери. Для пари випадкових величин X і Y такою додатковою парою літер буде (j, n_Y) , додатковою однією літерою – (l, L) , як у прикладі 4.6, а для величини X , що входить у пари, обсяг вибірки буде позначатися n_X . ▲

Зауваження 4.14. Обчислити вибіркові середні значення дискретних випадкових величин X і Y із прикладу 4.6 за формулою (4.5) неможливо, тому що вибірки не задані безпосередньо вибірковими значеннями $x_i, i = \overline{1, n_X}$ та $y_j, j = \overline{1, n_Y}$. Але таке задання вибірки не є обов'язкове. Дійсно, сума вибіркових значень x_i може бути заміненa сумою значень груп x_k із вибіркового ряду розподілу частот з урахуванням їх повторення:

$$\sum_{i=1}^n x_i = \sum_{k=1}^K n_k x_k,$$

а вираз (4.5) – таким, який також є визначенням вибіркового середнього значення випадкової величини X :

$$m = m_X = \frac{1}{n} \sum_{k=1}^K n_k x_k. \quad (4.6)$$

Аналогічний вираз можна одержати і для вибіркового середнього значення випадкової величини Y . ▲

Вибірковим розмахом дискретної випадкової величини X називається різниця між максимальним і мінімальним вибірковими значеннями:

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

Говорячи про вибірковий розмах дискретної випадкової величини X , варто мати на увазі, що вважаються відомими не тільки різниця між максимальним і мінімальним вибірковими значеннями, але й самі ці значення $x_{\min}, x_{\max}$. Тому можна визначити максимальне відхилення від вибіркового середнього значення в більшу ($x_{\max} - m_X$) і меншу ($x_{\min} - m_X$) сторони (з урахуванням знака відхилення).

Приклад 4.6 (продовження). $d_X = 14, = +7, x_{\min} - m_X = -7, d_Y = 8; y_{\max} - m_Y = +4, y_{\min} - m_Y = -4$. ●

Очевидно, що вибіркoве середнє значення і вибіркoвий розмах також недостатньо докладно описують дискретні випадкові величини.

Приклад 4.7. Для дискретних випадкових величин U і V , заданих вибірковими рядами розподілу частот, наведеними в табл. 4.13 і 4.14, є рівні чисельності груп повторюваних значень $K=L$, значення груп з однаковими номерами ($u_k = v_l$ при $k = l$), вибіркові середні $m_U = m_V = 6$, вибіркові відхилення $d_U = d_V = 10$ і відхилення в більшу і меншу сторони, однак ці випадкові величини мають якісно різні ряди розподілу.

Таблиця 4.13

u_k	1	6	11
n_k	10	20	10

Таблиця 4.14

	v_l	1	6	11	
	n_l	20	10	20	●

Як міру відмінності вибірових значень від вибірового середнього значення вводять **відхилення від вибірового середнього значення**. Їх можна ввести як для вибірових значень x_i , так і для груп вибірових значень x_k (однаково для дискретних і неперервних випадкових величин):

$$d_i = x_i - m, \quad i = \overline{1, n}. \quad \Bigg| \quad d_k = x_k - m, \quad k = \overline{1, K}.$$

Приклад 4.6 (продовження). Відхилення від вибірових середніх значень для дискретних випадкових величин X і Y , записані у вигляді рядів (по групах значень), наведені в табл. 4.15 і 4.16.

Таблиця 4.15

d_k	-3	-2	0	+2	+3
n_k	1	10	25	10	1

Таблиця 4.16

d_l	-1	0	0
n_l	20	10	20

 ●

Приклад 4.7 (продовження). Відхилення від вибірових середніх значень для дискретних випадкових величин U і V , записані у вигляді рядів (по групах значень), наведені в табл. 4.17 і 4.18.

Таблиця 4.17

d_k	-5	0	+5
n_k	10	20	10

Таблиця 4.18

	d_l	-5	0	+5
	n_l	20	10	20

 ●

Зауваження 4.16. Ряд відхилень від вибірових середніх значень разом із вибіровим середнім значенням є еквівалентні вибіровому ряду розподілу частот. ▲

Перейдемо до деякої узагальненої величини, що характеризує відхилення від вибірового середнього значення, наприклад шляхом осереднення відхилень для вибірових значень x_i або груп значень x_k з урахуванням частот груп n_k :

$$\begin{aligned} \frac{1}{n} \sum_{k=1}^K n_k d_k &= \frac{1}{n} \sum_{i=1}^n d_i = \frac{(x_1 - m) + (x_2 - m) + \dots + (x_n - m)}{n} = \\ &= \frac{x_1 + x_2 + \dots + x_n - n \cdot m}{n} = \frac{x_1 + x_2 + \dots + x_n}{n} - m = m - m = 0. \end{aligned}$$

Отже, середнє значення відхилень від вибіркового середнього значення дорівнює нулеві й не може бути узагальненою характеристикою відхилення вибірових значень від вибіркового середнього значення.

Замість осереднення відхилень застосовують середнє значення абсолютних відхилень від вибіркового середнього значення:

$$\frac{1}{n} \sum_{i=1}^n |d_i| = \frac{1}{n} \sum_{i=1}^n |x_i - m|$$

і середнє значення квадратів відхилень від вибіркового середнього значення:

$$s^2 = \frac{1}{n} \sum_{i=1}^n d_i^2 = \frac{1}{n} \sum_{i=1}^n (x_i - m)^2, \quad (4.7)$$

називане **вибірковою дисперсією** випадкової величини.

Найбільш часто застосовується вибіркова дисперсія, оскільки у визначенні цієї величини не використовуються абсолютні значення відхилень $|d_i|$, тобто вибіркова дисперсія як функція змінних $x_1, x_2, \dots, x_n$ є диференційована.

Поряд із вибірковою дисперсією випадкової величини зручно застосовувати додатний квадратний корінь із вибіркової дисперсії. Ця величина називається **вибірковим середньоквадратичним відхиленням**:

$$s = \sqrt{s^2} = \sqrt{\frac{1}{n} \sum_{i=1}^n (x_i - m)^2}. \quad (4.8)$$

Зауваження 4.17. Розмірність вибіркової дисперсії випадкової величини дорівнює квадратіві розмірності випадкової величини, тобто $[s_X^2] = [X]^2$. Наприклад, якщо значення випадкової величини вимірюються в кілограмах, метрах, відсотках і т.д., до обчисленої за результатами серії n дослідів вибіркової дисперсії (4.7) варто приписати праворуч розмірність кг^2 , м^2 , $\%^2$ і т.д. Розмірність же вибірових середньоквадратичних відхилень (4.8) дорівнює розмірності випадкової величини, і тому застосовувати цю величину дуже зручно. ▲

Вибіркове середнє значення і вибіркова дисперсія називаються **вибірковими числовими характеристиками випадкової величини**. Також відзначимо, що вибіркове середнє значення і вибіркова дисперсія не єдині вибірові числові характеристики, які ми будемо розглядати.

Вибіркові числові характеристики задають неповний опис випадкової величини, що не заміняє повного опису за вибіркою у вигляді вибіркового ряду розподілу (для дискретної випадкової величини) або інтервального ряду розподілу (для неперервної випадкової величини).

Вибіркові числові характеристики, які обчислюються за відповідними формулами для конкретної вибірки, дають цілком визначені числові значення. Для множини всіх можливих вибірок обсягу n з даної генеральної сукупності можна одержати різні значення числових характеристик, тобто на множині вибірових сукупностей заданого обсягу n визначені випадкові величини, які ґрунтуються на виразах для відповідних числових характеристик. Це означає, що для вибіркового середнього значення й вибіркової дисперсії для різних вибірок одержуються різні значення, що вказується в їх позначеннях у такий спосіб: $\tilde{m}$, $\tilde{s}^2$.

4.12. Властивості вибірових середнього значення та дисперсії

Властивості вибірових числових характеристик, а також їх доведення будемо подавати одночасно: ліворуч – для вибіркового середнього значення, праворуч – для вибіркової дисперсії.

1. Вибіркове середнє значення добутку сталої величини c і випадкової величини X дорівнює добуткові сталої величини c і вибіркового середнього значення випадкової величини X :

$$\begin{aligned} m_{cX} &= c m_X, \\ m_{cX} &= \frac{cx_1 + cx_2 + \dots + cx_n}{n} = \\ &= \frac{c(x_1 + x_2 + \dots + x_n)}{n} = c m_X. \end{aligned}$$

2. Вибіркове середнє значення суми сталої величини c і випадкової величини X дорівнює сумі сталої величини c і вибіркового середнього випадкової величини X :

$$m_{c+X} = c + m_X,$$

1. Вибіркова дисперсія добутку сталої величини c і випадкової величини X дорівнює добуткові квадрата сталої величини c і вибіркової дисперсії випадкової величини X :

$$\begin{aligned} s_{cX}^2 &= c^2 s_X^2, \\ s_{cX}^2 &= \frac{1}{n} \left\{ (cx_1 - cm_X)^2 + \dots \right. \\ &\quad \left. \dots + (cx_n - cm_X)^2 \right\} = \\ &= \frac{1}{n} \left\{ c^2 (x_1 - m_X)^2 + \dots \right. \\ &\quad \left. \dots + c^2 (x_n - m_X)^2 \right\} = \\ &= c^2 \frac{1}{n} \left\{ (x_1 - m_X)^2 + \dots \right. \\ &\quad \left. \dots + (x_n - m_X)^2 \right\} = c^2 s_X^2. \end{aligned}$$

2. Вибіркова дисперсія суми сталої величини c і випадкової величини X дорівнює вибіровій дисперсії випадкової величини X :

$$s_{c+X}^2 = s_X^2,$$

$$\begin{aligned}
 m_{c+X} &= \frac{(c+x_1) + \dots + (c+x_n)}{n} = \\
 &= \frac{nc + (x_1 + \dots + x_n)}{n} = \\
 &= c + \frac{(x_1 + \dots + x_n)}{n} = c + m_X
 \end{aligned}$$

3. Для довільних сталих величин a і b :

$$m_{aX+b} = a \cdot m_X + b.$$

4. Вибіркове середнє значення сталої величини c дорівнює значенню сталої величини:

$$m_c = c.$$

$$\begin{aligned}
 s_{c+X}^2 &= \frac{1}{n} \left\{ ((c+x_1) - (c+m_X))^2 + \dots \right. \\
 &\quad \left. \dots + ((c+x_n) - (c+m_X))^2 \right\} = \\
 &= \frac{1}{n} \left\{ (x_1 - m_X)^2 + \dots \right. \\
 &\quad \left. \dots + (x_n - m_X)^2 \right\} = s_X^2.
 \end{aligned}$$

3. Для довільних сталих величин a і b :

$$s_{aX+b}^2 = a^2 s_X^2.$$

4. Вибіркова дисперсія сталої величини c дорівнює нулеві:

$$s_c^2 = 0.$$

5. Випадкова величина X з вибірковою дисперсією, яка дорівнює нулеві, є сталою величиною.

$$\begin{aligned}
 s_X^2 &= \frac{1}{n} \left\{ (x_1 - m_X)^2 + \dots \right. \\
 &\quad \left. \dots + (x_n - m_X)^2 \right\} = 0 \Rightarrow \\
 (x_1 - m_X)^2 + \dots + (x_n - m_X)^2 &= 0 \\
 \Rightarrow \begin{bmatrix} d_1 = x_1 - m_X = 0, \\ d_2 = x_2 - m_X = 0, \\ \dots \\ d_n = x_n - m_X = 0; \end{bmatrix} &\Rightarrow \\
 \Rightarrow \begin{bmatrix} x_1 = m_X \\ x_2 = m_X \\ \dots \\ x_n = m_X \end{bmatrix} &\Rightarrow X \equiv m_X = \text{const.}
 \end{aligned}$$

4.13. Генеральні числові характеристики

Якщо поняття числових характеристик вибірки поширити на всю генеральну сукупність, тобто у виразах для m (4.2), (4.6) і s^2 (4.7) формально перейти до границі $n \rightarrow N(\infty)$, то одержані значення будуть називатися (**генеральним**) **середнім значенням** (або **математичним сподіванням**) μ і (**генеральною**) **дисперсією** σ^2

випадкової величини. Генеральне середнє значення й генеральна дисперсія називаються **(генеральними) числовими характеристиками випадкової величини** й задають неповний опис генеральної сукупності за ознакою (випадковою величиною) X , а повний опис задається рядом розподілу для дискретної випадкової величини та функцією густини для неперервної випадкової величини або функцією розподілу для дискретної та неперервної випадкових величин.

Зауваження 4.18. Іноді генеральні й вибіркові числові характеристики випадкової величини задають повний опис випадкової величини – як дискретної, так і неперервної. Це відбувається тоді, коли закон розподілу випадкової величини має який-небудь спеціальний вид. ▲

4.14. Зв'язок вибіркових і генеральних числових характеристик

Вибіркові числові характеристики, що є випадковими величинами, називаються **статистиками**, а (генеральні) числові характеристики, які є визначеними числовими значеннями (узагалі кажучи, невідомими), називаються **параметрами** випадкової величини (ознаки) X . Статистики, обчислені за конкретною вибіркою, дають наближені значення відповідних їм параметрів ($m \approx \mu$, $s^2 \approx \sigma^2$). Заміна невідомого значення параметра (генеральної числової характеристики) відомим значенням статистики (вибіркової числової характеристики) називається **точковою оцінкою параметра**. Походження терміна пов'язане з тим, що обчислена за даною вибіркою статистика на відповідній числовій осі може бути позначена точкою, розташованою на невідомій відстані від точки, що відповідає значенню параметра. Причому, не відомо, збільшує чи зменшує статистика значення параметра, тобто чи розташована точка, яка відповідає статистиці, ліворуч або праворуч точки, що відповідає параметрові. Недолік точкової оцінки компенсується **інтервальною оцінкою параметра**, тобто визначенням інтервалу, у якому з **надійною ймовірністю** (або **надійністю**, яка позначається β , $\beta \approx 1$) знаходиться невідоме значення параметра.

Параметри звичайно позначають грецькими літерами, а статистики – латинськими. Параметрами є, наприклад, математичне сподівання μ і дисперсія σ^2 .

4.15. Властивості генеральних середнього значення та дисперсії

Властивості генерального середнього значення й генеральної дисперсії можуть бути виведені з властивостей вибіркового середнього значення та властивостей вибіркової дисперсії (див. розд. 4.13).

Дійсно, властивості вибіркового середнього значення й вибіркової дисперсії не залежать від обсягу вибірки n і тому можуть бути поширені на вибірку, що збігається з генеральною сукупністю, тобто для випадку, коли $N = n$.

Властивості числових характеристик подамо одночасно: ліворуч – для генерального середнього, праворуч – для генеральної дисперсії:

1. Генеральне середнє значення добутку сталої величини c і випадкової величини X дорівнює добуткові сталої величини c і генерального середнього значення випадкової величини X :

$$\mu_{cX} = c \mu_X.$$

2. Генеральне середнє значення суми сталої величини c і випадкової величини X дорівнює сумі сталої величини і генерального середнього X :

$$\mu_{c+X} = c + \mu_X.$$

3. Для довільних сталих величин a і b :

$$\mu_{aX+b} = a \mu_X + b.$$

4. Генеральне середнє значення сталої величини c дорівнює значенню цієї сталої величини:

$$\mu_c = c.$$

1. Генеральна дисперсія добутку сталої величини c і випадкової величини X дорівнює добуткові квадрата сталої величини c і генеральної дисперсії випадкової величини X :

$$\sigma_{cX}^2 = c^2 \sigma_X^2.$$

2. Генеральна дисперсія суми сталої величини c і випадкової величини X дорівнює генеральній дисперсії випадкової величини X :

$$\sigma_{c+X}^2 = \sigma_X^2.$$

3. Для довільних сталих величин a і b :

$$\sigma_{aX+b}^2 = a^2 \sigma_X^2.$$

4. Генеральна дисперсія сталої величини c дорівнює нулеві:

$$\sigma_c^2 = 0.$$

5. Випадкова величина X з генеральною дисперсією, що дорівнює нулю, є сталою величиною.

4.16. Опис випадкової величини

Розглянуті раніше різні види повного й неповного опису випадкових величин зведені в табл. 4.21. Покажемо, що з повного опису випадкової величини можна одержати наближене, але не навпаки.

Визначення числових характеристик дискретної випадкової величини за її повним описом

У розд. 4.11 було показано, що вибіркове середнє значення подається як через вибірккові значення x_i , $i = \overline{1, n}$, так і через значення груп x_k , $k = \overline{1, K}$. У виразі (4.6) можна явно виділити відносні частоти, які за виконання умови стабілізації (див. розд. 3.2) можна замінити статистичними ймовірностями:

Таблиця. 4.19

Випадкова величина	Опис випадкової величини			
	Повний		Ненормований	
Дискретна	У генеральній сукупності (точний)	У вибірковій сукупності (наближений)	У генеральній сукупності (точний)	У вибірковій сукупності (наближений)
	(генеральний) ряд розподілу	вибірковий ряд розподілу	а) генеральне середнє значення (математичне сподівання) μ ; б) генеральна дисперсія σ^2	а) вибіркве середнє значення m ; б) вибіркво дисперсія n -зсунена $\hat{s}^2$.
	а) (генеральна) функція густини $f(x)$; б) (генеральна) функція розподілу $F(x)$	а) вибіркво функція густини $f^*(x)$; б) вибіркво функція розподілу $F^*(x)$		
Неперервна				

$$m = m_X = \sum_{k=1}^K \frac{n_k}{n} x_k = \sum_{k=1}^K p_k^* x_k . \quad (4.9)$$

Якщо в вибірку попадає вся генеральна сукупність ($n \rightarrow N$), то статистичні ймовірності замінюються на ймовірності, а вибіркове середнє значення – на генеральне середнє значення:

$$\mu = \mu_X = M[X] = \sum_{k=1}^K p_k x_k . \quad (4.10)$$

Відзначимо, що у разі обчислення вибіркового середнього значення за формулою (4.9) і генерального середнього значення за формулою (4.10) враховують вибірковий ряд розподілу відносних частот і (генеральний) ряд розподілу. Для генерального середнього значення це виявляється у використанні позначення $M[X]$, у якому підкреслюється, що для обчислення генерального середнього значення у квадратні дужки $[]$ заноситься не яке-небудь одне значення дискретної випадкової величини, а вся множина її значень разом з ймовірностями їх появи, тобто (генеральний) ряд розподілу.

Іншою вибірковою числовою характеристикою є вибіркова дисперсія. Очевидно, що у визначенні вибіркової дисперсії (4.7) сума квадратів відхилень від вибіркового середнього значення замінюється сумою квадратів відхилень значень груп d_k^2 від вибіркового середнього значення, які помножені на частоти n_k значень груп x_k :

$$\sum_{i=1}^n d_i^2 = \sum_{k=1}^K n_k d_k^2 .$$

За виконання умови стабілізації відносних частот

$$s^2 = \sum_{k=1}^n \frac{n_k}{n} d_k^2 = \sum_k p_k^* d_k^2 . \quad (4.11)$$

Границя (4.11) при $n \rightarrow N$ дає такий вираз для (генеральної) дисперсії:

$$\sigma^2 = \sigma_X^2 = D[X] = \sum_{k=1}^K p_k (x_k - \mu_X)^2 . \quad (4.12)$$

Позначення $D[X]$ указує на те, що при обчисленні (генеральної) дисперсії використовується ряд розподілу дискретної випадкової величини.

Визначення числових характеристик неперервної випадкової величини за її повним описом

Для того щоб одержати вираз генерального середнього значення через функцію густини (тобто через закон розподілу неперервної випадкової величини), так само, як і у випадку дискретної випадкової величини, застосуємо визначення вибіркового середнього значення (4.5), (4.6). Розподілені по інтервалах вибіркові значення втрачають свою індивідуальність, тому замість них будемо використовувати які-небудь характерні інтервальні значення. Наприклад, для інтервалу $[a_k, a_{k+1})$ з номером k характерним візьмемо середину інтервалу:

$$x_k^* = \frac{1}{2}(a_k + a_{k+1}) = a_k + \frac{1}{2}\Delta a.$$

Потім вираз для суми вибірових значень замінимо сумою характерних інтервальних значень x_k^* , помножених на інтервальні частоти n_k :

$$\sum_{i=1}^n x_i \approx \sum_{k=1}^K n_k x_k^*. \quad (4.13)$$

Зверніть увагу на те, що ця заміна наближена, тому вираз для вибіркового середнього є наближений:

$$m_X \approx \sum_{k=1}^K \frac{n_k}{n} x_k^*. \quad (4.14)$$

Звідси вже легко одержати остаточний вираз для генерального середнього значення. Для цього достатньо перейти до границі при $\Delta a \rightarrow 0$ ¹ ($K \rightarrow \infty$) і $n \rightarrow N$. По-перше, за умови стабілізації відносних частот замінимо їх статистичними ймовірностями:

¹ До речі, саме умова $\Delta a \rightarrow 0$ дозволяє як характерне інтервальне значення взяти будь-яке значення з відповідного інтервалу (не обов'язково його середину, адже межі інтервалу зближуються), хоч очевидно, що при кінцевій ширині інтервалу вибору середини інтервалу як x_k^* слід віддавати перевагу.

$$m_X \approx \sum_{k=1}^K \frac{n_k}{n} x_k^* = \sum_{k=1}^K p_k^* x_k^* . \quad (4.15)$$

По-друге, виразимо у формулі (4.15) p_k^* через добуток вибіркової функції густини f_k^* і ширини інтервалу Δa (див. розд. 4.10):

$$m_X \approx \sum_{k=1}^K f_k^* \Delta a x_k^* .$$

Перед тим як зробити формальний перехід $\Delta a \rightarrow 0$ і $n \rightarrow N$:

- 1) замінимо позначення ширини інтервалу Δa на Δx ;
- 2) змінимо порядок співмножників:

$$m_X \approx \sum_{k=1}^K f_k^* x_k^* \Delta x . \quad (4.16)$$

Вираз (4.16) – це інтегральна сума для функції $x f(x)$. Отже, після граничного переходу (за змістом границя існує) одержимо шуканий вираз для генерального середнього значення (математичного сподівання) неперервної випадкової величини:

$$\mu = \mu_X = M[X] = \int_{-\infty}^{+\infty} x f(x) dx . \quad (4.17)$$

Нижня та верхня межі інтегрування у виразі (4.17) згідно з зауваженням 4.12 записані відповідно як $-\infty$ і $+\infty$. Безпосередньо обчислювати невласний інтеграл (4.17) не потрібно. За його допомогою покажемо, як із повного опису одержати одну з генеральних числових характеристик, а саме генеральне середнє значення. Для того щоб це підкреслити, використовують позначення $M[X]$. Як і для дискретної випадкової величини, у цьому записі мається на увазі, що у квадратні дужки “заноситься” і закон розподілу (у вигляді (генеральної) функції густини $f(x)$), і множина значень неперервної випадкової величини (як область визначення $f(x)$), тобто неперервна випадкова величина “цілком”. Тому й використовується велика літера X у квадратних дужках. Вона відіграє зовсім іншу роль, ніж мала літера x у позначенні будь-якої функції $y(x)$. У функції одному числу (аргументу x) ставиться у відповідність інше число – значення функції $y(x)$. У виразі (4.17) “усій” неперервній випадковій величині ставиться у відповідність одне число – генеральне середнє значення.

Аналогічно вводиться вираз (генеральної) дисперсії через (генеральну) функцію густини. З огляду на повну подібність із виведенням генерального середнього значення далі наводимо тільки ланцюжок перетворень без коментарів:

$$\begin{aligned} s^2 &= \frac{1}{n} \sum_{i=1}^n (x_i - m_X)^2 \approx \frac{1}{n} \sum_{k=1}^K n_k (x_k^* - m_X)^2 = \\ &= \sum_{k=1}^K p_k^* (x_k^* - m_X)^2 = \sum_{k=1}^K f_k^* \Delta a (x_k^* - m_X)^2; \end{aligned} \quad (4.18)$$

$$\sigma^2 = \sigma_X^2 = D[X] = \lim_{\Delta x \rightarrow 0} \sum_{k=1}^K (x_k^* - m_X)^2 f_k^* \Delta x = \int_{-\infty}^{+\infty} (x - \mu_X)^2 f(x) dx. \quad (4.19)$$

Зауваження 4.19. Уведена раніше вибіркова дисперсія “добре” “працює” тільки за великих обсягів вибірки, а у випадку малих обсягів – систематично знижує (зсуває) оцінку (генеральної) дисперсії. Тому s^2 (для дискретної та неперервної випадкових величин) часто називають **зсуненою вибірковою дисперсією**. При виконанні первинної обробки вибірових даних в обчисленнях переважно застосовується незсунена вибіркова дисперсія (зверніть увагу на зміну позначення):

$$\hat{s}^2 = \hat{s}_X^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - m_X)^2. \quad (4.20)$$

Докладніше про це буде йти мова далі під час розгляду задач оцінки параметрів.

4.17. Квантилі неперервної випадкової величини

Квантилем x_q неперервної випадкової величини X називається її значення, що задовольняє таку умову:

$$P(\{X < x_q\}) = q, \quad 0 \leq q \leq 1,$$

де q – індекс квантиля.

Поняття квантиля містить два компоненти:

- 1) певне значення x неперервної випадкової величини X ;
- 2) індекс q , що відповідає значенню x_q неперервної випадкової величини й дорівнює значенню функції розподілу, тобто $q = F(x_q)$ (рис. 4.22). Якщо згадати ви- значення функції розподілу (див. розд. 4.15), то про індекс q квантиля можна гово-

рити як про ймовірність набуття неперервною випадковою величиною X значень, менших значення квантиля x_q . А з огляду на властивість 4 функції густини індекс виявляється таким, що дорівнює площі під кривою $f(x)$ ліворуч від x_q (рис. 4.23). Те саме можна подати в такому вигляді:

$$q = F(x_q) = \int_{-\infty}^{x_q} f(x) dx.$$

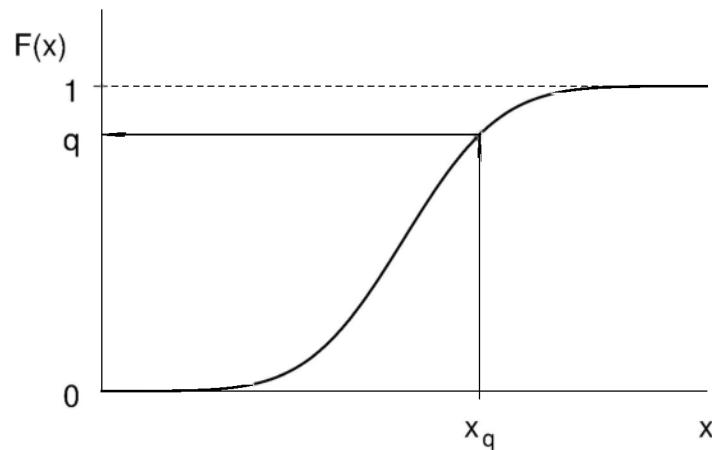


Рис. 4.22. Схема визначення квантиля за допомогою функції розподілу

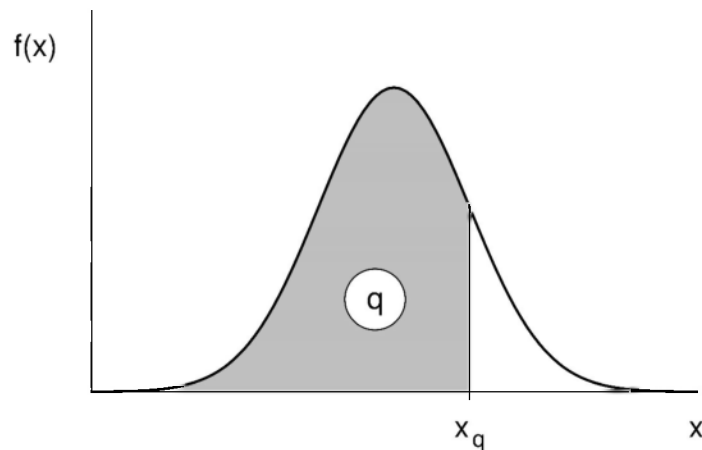


Рис. 4.23. Схема визначення квантиля за допомогою функції густини

Оскільки функція розподілу $F(x)$ монотонна (див. властивість 3 у розд. 4.14), то довільному значенню x неперервної випадкової величини X відповідає одне значення індексу $q = F(x)$ і, навпаки, довільному числу $0 \leq q \leq 1$ відповідає одне значення x_q неперервної випадкової величини X , що задовольняє визначення квантиля. Цей висновок впливає, зокрема, з графіка $F(x)$. Це можна трактувати так, що будь-

яке значення неперервної випадкової величини є квантилем із відповідним індексом.

Властивості квантилів

1. Для довільних квантилів x_{q_1} і x_{q_2} (рис. 4.24, 4.25)

$$\{x_{q_1} < x_{q_2}\} \Leftrightarrow \{q_1 < q_2\}.$$

2. Ймовірність набуття неперервною випадковою величиною значень між двома квантилями x_{q_1} і x_{q_2} дорівнює різниці їх індексів (рис. 4.24, 4.25)

$$P(\{x_{q_1} \leq X < x_{q_2}\}) = q_2 - q_1.$$

3. Альтернативне визначення квантиля (див. рис. 4.22, 4.23)

$$P(\{X > x_q\}) = 1 - q.$$

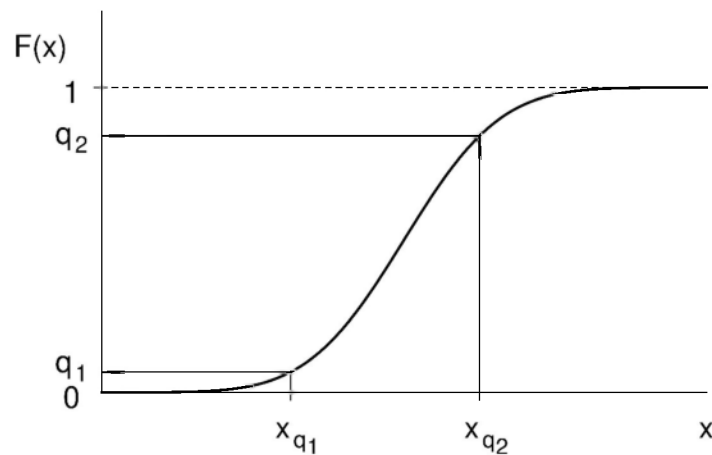


Рис. 4.24. Обґрунтування властивостей квантилів за допомогою функції розподілу

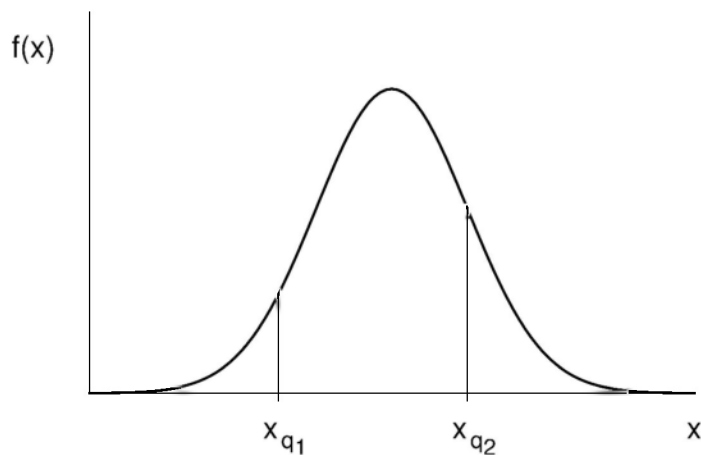


Рис. 4.25. Обґрунтування властивостей квантилів за допомогою функції густини

Для стандартних і похідних від них розподілів неперервної випадкової величини складені таблиці квантилів, тобто пари значень (q, x_q) . Найчастіше таблиці складають, виходячи з рівномірного розбиття відрізка $[0,1]$ за q , а не з розбиття множини значень неперервної випадкової величини X . Для деяких видів розбиття використовуються спеціальні назви. Наприклад “квартилі” для розбиття на чотири частини (рис. 4.26), “децилі” та “центилі” – для розбиття на десять і сто частин відповідно.

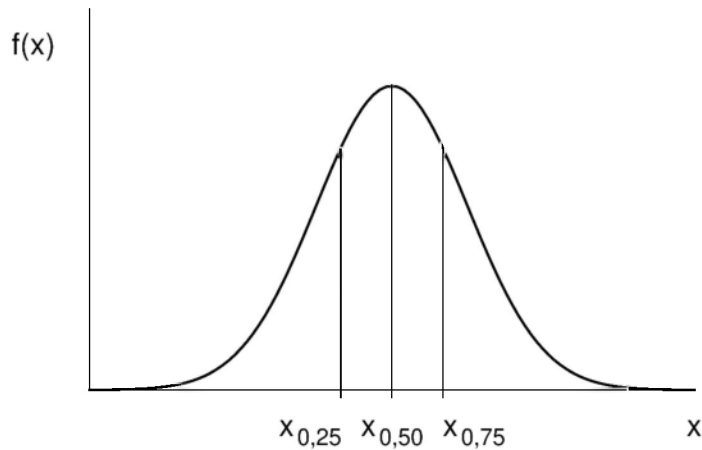


Рис. 4.26. Схема визначення квартилів

Оскільки довільна функція може бути задана аналітично (формулою), графічно або у вигляді таблиці, то за великої кількості квантилів (дрібного розбиття за q), таблиця квантилів досить докладно описує функцію розподілу $F(x)$ (табл. 4.19).

Квантили найбільш часто використовуються при перевірці статистичних гіпотез, причому достатньо знати квантили при $q \approx 0$ і $q \approx 1$.

Таблиця 4.20

k	q	x_q
1	q_1	x_{q_1}
2	q_2	x_{q_2}
...	...	...
k	q_k	x_{q_k}
...	...	...

При використанні квантилів, як правило, доводиться розв’язувати дві **задачі** – **пряму й обернену**. У **прямій задачі** за значенням x неперервної випадкової величини X необхідно визначити його індекс q як квантиля, в **оберненій** – навпаки, за індексом q необхідно знайти значення x неперервної випадкової величини X , що відповідає даному індексові. При обробці даних і перевірці статистичних гіпотез за таблицями квантилів обидві задачі розв’язуються наближено. Якщо в колонці значень q чи x_q відсутні потрібні значення, то застосовується лінійна інтерполяція.

4.18. Побудова моделі випадкової величини за її описом

У попередніх розділах були введені поняття прихованого механізму і множини значень випадкової величини.

По-перше, з аналізу даних серії дослідів із дискретними або неперервними випадковими величинами можна наближено встановити множину їх значень. По-друге, при вивченні закономірностей появи тих чи інших значень випадкової величини були введені такі поняття, як вибіркового і генерального ряди розподілу, вибіркова і генеральна функції густини і розподілу. За їх допомогою можна відтворити механічні моделі прихованих механізмів випадкових величин.

Почнемо, як звичайно, із дискретної випадкової величини, а потім перейдемо до неперервної випадкової величини.

Механічна модель дискретної випадкової величини

Найпростішою є побудова для вибіркового ряду розподілу частот. Візьмемо коробку й заповнимо її кульками однакового розміру, на яких написані значення груп x_k (кожне таке значення буде визначати вид, або “сорт” кульок), тобто кількість кульок зі значенням $x_{k=1}$ дорівнює $n_{k=1}$ і т.д. Якщо помістити кульки в коробку, добре перемішати й закрити, а потім “випадково”, без спроб тенденційного добору (наскільки це можливо), витягати їх із поверненням, то буде одержана механічна модель дискретної випадкової величини для вибіркового ряду – частини генеральної сукупності, на якій визначена дискретна випадкова величина. При кожному витягненні записується “вибіркоче” значення дискретної випадкової величини (це моделює процес одержання вибіркового значення випадкової величини в досліді шляхом вимірювання), однак кількість дослідів може бути довільною (нагадаємо, що для побудови вибіркового ряду розподілу частот у розділі 4.11 було використано n дослідів). Оскільки кількість кульок кожного виду відома, то в даних дослідів виконуються всі умови для обчислення класичної ймовірності подій $\{X = x_k\}$. Проте, продовжуючи витягнення, досить довго можна спостерігати стабілізацію відносних частот і визначити статистичні ймовірності для значень груп x_k . Зрозуміло, що повинна виконуватися наближена рівність між ними та відповідними статистичними ймовірностями (див. табл. 4.6). У дослідях із механічною моделлю останні уже виступають як класичні ймовірності.

Якщо заданий точний повний опис дискретної випадкової величини, тобто її ряд розподілу, побудова механічної моделі проводиться в такий спосіб. Спочатку задається кількість кульок n , що будуть поміщені в коробку. Потім для кожного значення груп x_k (значення дискретної випадкової величини) визначається відповідна кількість кульок:

$$n_k \approx n p_k.$$

Далі все відбувається так само, як і у випадку побудови механічної моделі за вибіркою. Зрозуміло, що уточнення моделі відбувається за умови $n \rightarrow N$.

Потрібно також враховувати, що наші уявлення про випадковість появи тих або

інших значень x_k у досліді можуть не зовсім точно відтворювати витягнення кульок виду x_k з коробки.

Механічна модель неперервної випадкової величини

Якщо відомий вибірковий інтервальний ряд розподілу частот, то потрібно лише визначити характерні значення інтервалів x_k^* (наприклад, як у розд. 4.16), а далі виконувати дії, аналогічні тим, що мають місце при побудові моделі дискретної випадкової величини за вибіркоvim рядом розподілу частот.

Побудуємо модель неперервної випадкової величини за відомим описом неперервної випадкової величини у вигляді (генеральної) функції густини. Уведемо K інтервалів шириною Δa , що розбивають область визначення функції $f(x)$ (для неперервної випадкової величини це і є множина її значень!). Якщо $f(x)$ визначена на скінченному проміжку $[x_{\min}, x_{\max}]$, за межами якого вона дорівнює нулеві (про такі функції говорять, що вони мають компактний носій), то кількість інтервалів K є скінченна. Насправді тільки такі ситуації, у яких неперервна випадкова величина X має обмежену знизу і зверху множину значень, і мають місце в реальних задачах. Але в ряді дослідів надійний опис неперервної випадкової величини дається за допомогою якої-небудь стандартної функції густини, визначеної на всій числовій осі $(-\infty, +\infty)$. Тому кількість інтервалів K не обмежена, тобто $K=\infty$. Розв'язання цієї задачі ґрунтується на тому, що ймовірність набуття неперервною випадковою величиною з такою функцією густини значень, які знаходяться далеко від початку координат $x = 0$ і йдуть у бік $-\infty$ і $+\infty$, дуже швидко зменшується (гілки графіка функції густини, що йдуть у бік $-\infty$ і $+\infty$, швидко наближаються до осі x). Тому дуже великі за модулем значення неперервної випадкової величини виявляються малоймовірними. Практично це означає таке. Якщо обчислити площі S_k під графіком функції густини в межах кожного інтервалу k , то будуть визначені ймовірності $p_k = S_k$ набуття неперервною випадковою величиною X значень із цих інтервалів (див. властивість 3 функції густини в розд. 4.11). Тепер задамо загальну кількість кульок n , які будуть використовуватись у побудові механічної моделі, а кількість кульок n_k , що характеризують інтервал k , визначимо з наближеної рівності

$$n_k \approx n p_k.$$

Якщо ймовірності p_k для деяких інтервалів дуже малі, то вважатимемо, що $n_k \approx 0$, тобто кульки, що представляють даний інтервал, будуть відсутні. На цьому побудова механічної моделі неперервної випадкової величини за її точним повним описом завершується.

Зверніть увагу на те, що точність моделі можна поліпшити двома способами: по-перше, збільшуючи n , по-друге, зменшуючи ширину інтервалу Δa , тобто збільшуючи кількість інтервалів K .

5. Стандартні розподіли дискретних випадкових величин

5.1. Біномний розподіл

Дискретна випадкова величина X має біномний розподіл (говорять також, що дискретна випадкова величина X розподілена за біномним законом), якщо її можливі значення $x \in X = \{0, 1, 2, \dots, m, \dots, n\}$, а відповідні їм ймовірності дорівнюють

$$b(m; n, p) = P(\{X = k\}) = C_n^m p^m q^{n-m}, \quad (5.1)$$

де $0 < p < 1$, $q = 1 - p$, $m = 0, 1, 2, \dots, n$.

Біномний розподіл є двохпараметричний, тому що залежить від двох параметрів m і p . Формула для $b(m; n, p)$ була одержана Я. Бернуллі¹ й названа його ім'ям.

На практиці біномний розподіл виникає за таких умов, що складають **схему Бернуллі**:

- 1) проводиться n дослідів;
- 2) у кожному з дослідів з ймовірністю p настає подія A , причому ймовірність p не змінюється від досліду до досліду.

Покажемо, що схема Бернуллі приводить до біномного розподілу (5.1). Для цього розглянемо випадкову величину $X = \{\text{кількість появ події } A \text{ в схемі Бернуллі}\}$. Подія $B = \{X = m\}$ розпадається на такі варіанти:

$$B = B_1 + B_2 + \dots + B_s,$$

у кожному з яких m разів настає подія A і $(n - m)$ разів – раз подія $\bar{A}$, але порядок появи подій A і $\bar{A}$ різний. Наприклад, один із варіантів є такий:

$$B_1 = \left\{ \underbrace{A, A, \dots, A}_m, \underbrace{\bar{A}, \bar{A}, \dots, \bar{A}}_{n-m} \right\}.$$

Будь-який інший варіант $B_2, B_3, \dots, B_s$ можна одержати з варіанта B_1 зміною порядку настання подій A і $\bar{A}$.

За правилом множення ймовірностей для довільного числа незалежних подій (див. розд. 4.3)

$$P(B_1) = p^m (1 - p)^{n-m},$$

¹ Бернуллі Якоб (1654–1705) – професор математики Базельського університету (Швейцарія).

або, при позначенні $q = 1 - p$, $P(B_1) = p^m q^{n-m}$. Очевидно, що таку ж ймовірність буде мати будь-який інший варіант події B . Кількість таких варіантів дорівнює кількості способів, якими можна з n дослідів “вибрати” m таких дослідів, у яких настає подія A (див. розд. 2.3.), тобто $s = C_n^m$. Звідси за правилом додавання ймовірностей для довільної кількості несумісних подій (див. розд. 4.2.) одержуємо:

$$P(B) = C_n^m p^m q^{n-m},$$

тобто випадкова величина X має біномний розподіл.

Ряд розподілу такої випадкової величини показаний у вигляді табл. 5.1.

Таблиця 5.1

k	1	2	...	$m+1$	...	$n+1$
x_k	0	1	...	M	...	n
p_k	$C_n^0 p^0 q^{n-0} =$ $= q^n$	$C_n^1 p^1 q^{n-1}$	...	$C_n^m p^m q^{n-m}$	...	$C_n^n p^n q^{n-n} =$ $= p^n$

Можна також показати, що

$$\sum_{m=0}^n b(m; n, p) = \sum_{m=0}^n C_n^m p^m q^{n-m} = (p + q)^n = 1, \quad (5.2)$$

тобто розподіл ймовірностей відповідає коефіцієнтам розкладу **бінома Ньютона**.

Визначимо числові характеристики (параметри) біномного розподілу:

$$\begin{aligned} \mu = M[X] &= \sum_{m=0}^n m \cdot b(m; n, p) = \sum_{m=0}^n m \cdot C_n^m p^m q^{n-m} = \\ &= \sum_{m=0}^n m \cdot \frac{n!}{m!(n-m)!} p^m q^{n-m} = \sum_{m=0}^n np \frac{(n-1)!}{(m-1)!(n-m)!} p^{m-1} q^{n-m} = \quad (5.3) \\ &= np \sum_{m=0}^{n-1} C_{n-1}^m p^m q^{n-1-m} = np, \\ \sigma^2 = D[X] &= \sum_{m=0}^n m^2 \cdot b(m; n, p) - (np)^2 = \end{aligned}$$

$$\begin{aligned}
&= \sum_{m=0}^n m(m-1)C_n^m p^m q^{n-m} + \sum_{m=0}^n m \cdot C_n^m p^m q^{n-m} - (np)^2 = \\
&= \sum_{m=0}^n n(n-1)p^2 C_{n-2}^{m-2} p^{m-2} q^{(n-2)-(m-2)} + np - (np)^2 = \quad (5.4) \\
&= n(n-1)p^2 + np - (np)^2 = np(1-p) = npq.
\end{aligned}$$

Розглянемо біномний розподіл для значень параметрів $p = 0,1; 0,3; 0,5; 0,7; 0,9$ і $n = 8$. Числові характеристики такого розподілу наведені в табл. 5.2.

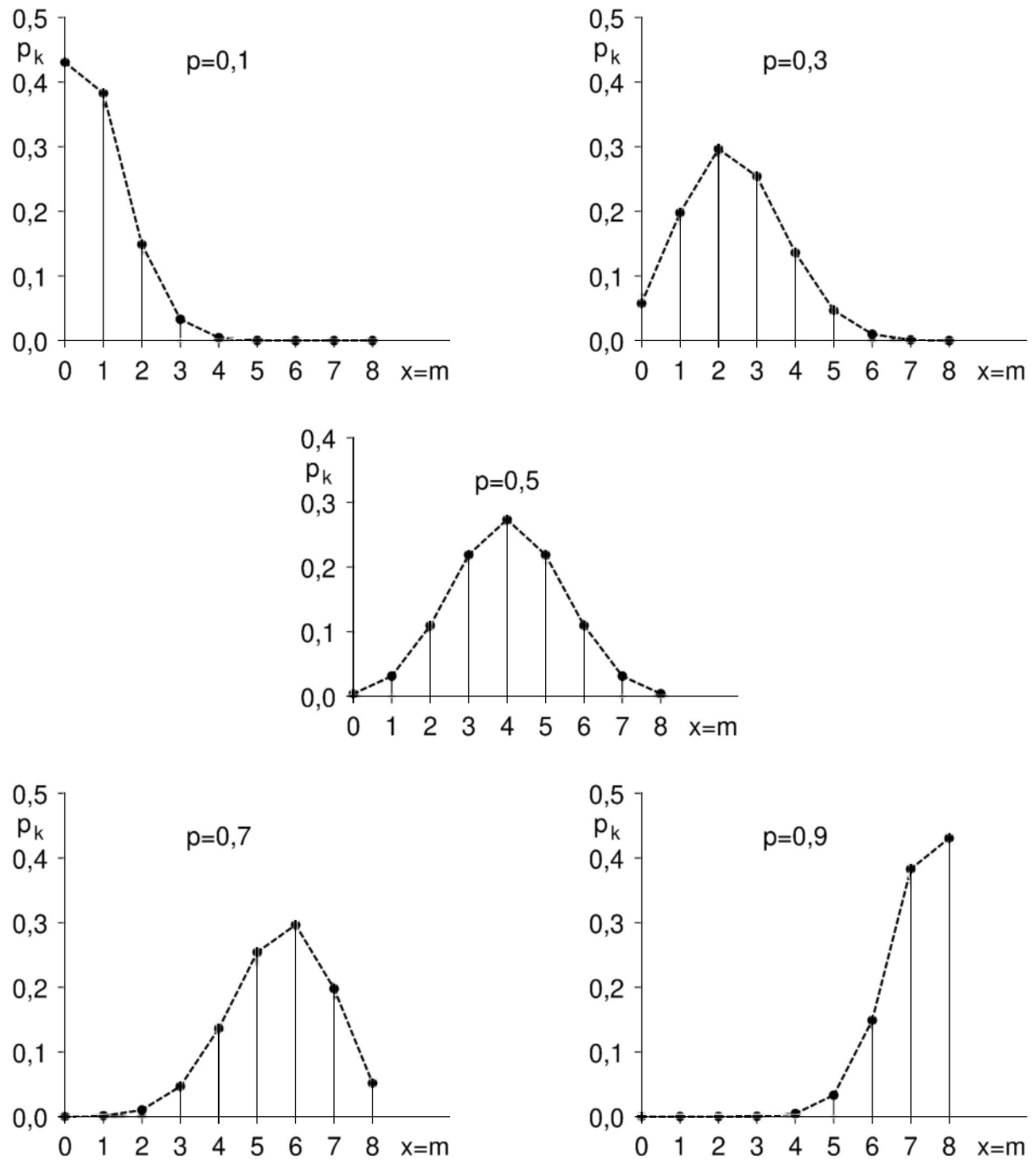
Таблиця 5.2

p	μ	σ
0,1	0,8	0,85
0,3	2,4	1,30
0,5	4,0	1,41
0,7	5,6	1,30
0,9	7,2	0,85

Відповідні ряди розподілів наведені в табл. 5.3, а графічне зображення рядів подане на рис. 5.1.

Таблиця 5.3

k	1	2	3	4	5	6	7	8	9
x_k	0	1	2	3	4	5	6	7	8
P_k ($p = 0,1$)	0,4305	0,3826	0,1489	0,0331	0,0046	0,0004	0,0000	0,0000	0,0000
P_k ($p = 0,3$)	0,0576	0,1977	0,2965	0,2541	0,1361	0,0467	0,0100	0,0012	0,0001
P_k ($p = 0,5$)	0,0039	0,0312	0,1094	0,2188	0,2734	0,2188	0,1094	0,0312	0,0039
P_k ($p = 0,7$)	0,0001	0,0012	0,0100	0,0467	0,1361	0,2541	0,2965	0,1977	0,0516
P_k ($p = 0,9$)	0,0000	0,0000	0,0000	0,0004	0,0046	0,0331	0,1489	0,3826	0,4305

Рис. 5.1. Біномний розподіл залежно від параметра p

Приклад 5.1. (Фішер, 1958). У 53 079 родин, що мають по 8 дітей, спостерігався розподіл за кількістю хлопчиків m ($m = \overline{0,8}$, $k = m+1$), наведений у табл. 5.4.

Таблиця 5.4

k	1	2	3	4	5	6	7	8	9
m	0	1	2	3	4	5	6	7	8
n_k	215	1 485	5 331	10 649	14 959	11 929	6 078	2 091	342

Якщо припустити, що розподіл дітей за статтю є біномний із параметром $p = 0,5$ ($q = 1 - p = 0,5$), то можна порівняти вибіркові дані у вигляді відносних частот n_k/n із теоретичними ймовірностями p_k , наведеними в табл. 5.5.

Таблиця 5.5

k	1	2	3	4	5	6	7	8	9
m	0	1	2	3	4	5	6	7	8
$\frac{n_k}{n}$	0,004051	0,02798	0,1004	0,2006	0,2818	0,2247	0,1145	0,03939	0,006443
p_k	0,003906	0,03125	0,1094	0,2188	0,2734	0,2188	0,1094	0,03125	0,003906

Графічне зображення дослідного і теоретичного (біномного) розподілів подане на рис. 5.2. ●

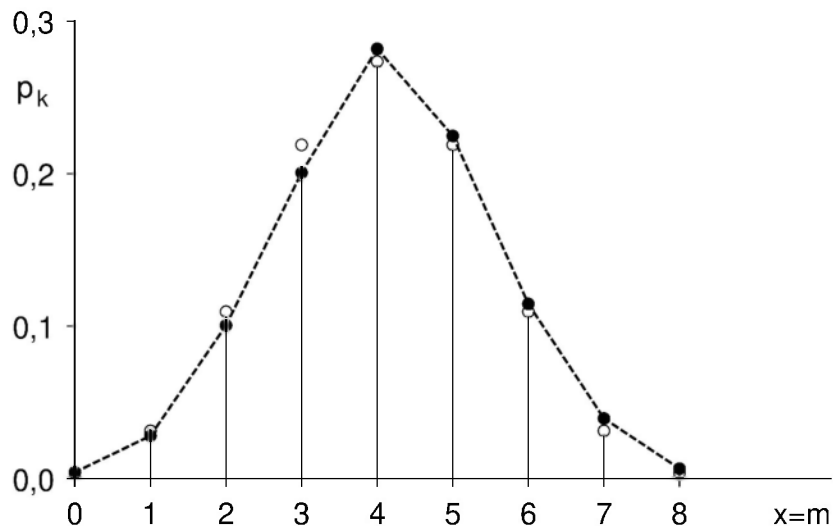


Рис. 5.2. Біномний розподіл (білі кружечки – теоретичний, чорні – емпіричний)

Приклад 5.2. (Лакін, 1982). Результати 20 160 підкидань чотирьох однакових монет наведені в табл. 5.6.

Таблиця 5.6

Комбінація		Частота комбінації	Відносна частота	Ймовірність комбінації
“орлів”	“решок”			
4	0	1181	0,0586	0,0625
3	1	4909	0,2435	0,2500
2	2	7583	0,3761	0,3700
1	3	5085	0,2522	0,2500
0	4	1402	0,0695	0,0625
Усього		20160	1,0000	1,0000

Для застосування біномного розподілу важливо забезпечити виконання умов схеми Бернуллі – сталість ймовірності p і незалежність випробувань, тобто тих умов, яких важко дотриматися на практиці.

Порушення незалежності дослідів може відбуватись у таких випадках:

1) якщо в досліді одержана відповідь “так” (наприклад, у діагностичному дослідженні виявлене невиконання діагностувального тесту), то приймається рішення “так” і послідовність експериментів обривається; якщо отримана відповідь “ні”, дослід триває доти, поки не буде вичерпана послідовність в n дослідів;

2) при опитуванні групи людей, коли, зокрема, дозволяється слухати відповіді інших.

5.2. Розподіл Пуассона

Дискретна випадкова величина X має **розподіл Пуассона**² (також говорять, що дискретна випадкова величина X розподілена за законом Пуассона), якщо її можливі значення $x \in X = \{0, 1, 2, \dots, m, \dots\}$, а відповідні їм ймовірності дорівнюють

$$P_{\lambda}(m) = P(\{X = m\}) = \frac{\lambda^m}{m!} e^{-\lambda}, \quad (5.5)$$

де $\lambda > 0$, $m = 0, 1, 2, \dots$.

Уперше розподіл Пуассона був описаний у його статті “Дослідження про ймовірність вироків у карних і цивільних справах”, опублікований у 1837 р.

Розглянемо умови, за яких виникає пуассонів розподіл. Насамперед покажемо, що він є граничний для біномного розподілу $b(m; n, p)$, коли кількість дослідів n необмежено збільшується ($n \rightarrow \infty$) і одночасно ймовірність p настання події A в одному досліді необмежено зменшується ($p \rightarrow 0$), але так, що їх добуток залишається сталим: $np = \lambda = \text{const}$.

Теорема Пуассона. Нехай p_n – ймовірність події A у кожному досліді серії з n дослідів схеми Бернуллі. Тоді якщо $\lambda_n = np_n \rightarrow \lambda$ ($0 < \lambda < \infty$) при $n \rightarrow \infty$, то біномний розподіл $b(m; n, p)$ при $n \rightarrow \infty$ збігається до розподілу Пуассона $P_{\lambda}(m)$.

Доведення. Подамо ймовірність у схемі Бернуллі у вигляді

$$\begin{aligned} b(m; n, p) &= C_n^m p_n^m q_n^{n-m} = \frac{n(n-1)\dots(n-m+1)}{m!} \left(\frac{\lambda_n}{n}\right)^m \left(1 - \frac{\lambda_n}{n}\right)^{n-m} = \\ &= \frac{\lambda_n^m}{m!} \cdot \frac{n}{n} \cdot \frac{n-1}{n} \cdot \dots \cdot \frac{n-m+1}{n} \left(1 - \frac{\lambda_n}{n}\right)^{-m} \left(1 - \frac{\lambda_n}{n}\right)^n \end{aligned}$$

² Пуассон (Poisson) Сімеон Дені (1781 – 1840) – французький фізик, математик, член Паризької академії наук (1812), іноземний член Петербурзької академії наук (1826).

і розглянемо граничний випадок $\lim_{n \rightarrow \infty} b(m; n, p)$:

$$\begin{aligned} & \lim_{n \rightarrow \infty} \frac{\lambda^m}{m!} \cdot \frac{n}{n} \cdot \frac{n-1}{n} \cdot \dots \cdot \frac{n-m+1}{n} \cdot \left(1 - \frac{\lambda_n}{n}\right)^{-m} \left(1 - \frac{\lambda_n}{n}\right)^n = \\ & = \frac{\lambda^m}{m!} \cdot \frac{\lim_{n \rightarrow \infty} \left(1 - \frac{\lambda_n}{n}\right)^n}{\lim_{n \rightarrow \infty} \left(1 - \frac{\lambda_n}{n}\right)^{-m}} = \frac{\frac{\lambda^m}{m!}}{\left[\lim_{n \rightarrow \infty} \left(1 - \frac{\lambda_n}{n}\right)^{-\frac{n}{\lambda_n}} \right]^{-\lambda_n}} = \frac{\lambda^m}{m!} e^{-\lambda} = P_\lambda(m). \end{aligned}$$

При доведенні використана “друга чудова границя” функції

$$\lim_{x \rightarrow \infty} \left(1 + \frac{1}{x}\right)^x = e. \blacksquare$$

Обчислимо числові характеристики (параметри) розподілу Пуассона:

$$\begin{aligned} \mu = M[X] &= \sum_{m=0}^{\infty} m \frac{\lambda^m}{m!} e^{-\lambda} = e^{-\lambda} \lambda \sum_{m=1}^{\infty} \frac{\lambda^{m-1}}{(m-1)!} = \\ &= \lambda e^{-\lambda} \sum_{m=1}^{\infty} \frac{\lambda^{m-1}}{(m-1)!} = \lambda e^{-\lambda} e^{\lambda} = \lambda; \end{aligned} \quad (5.6)$$

$$\begin{aligned} \sigma^2 = D[X] &= \sum_{m=0}^{\infty} m^2 \frac{\lambda^m}{m!} e^{-\lambda} - \lambda^2 = \lambda \sum_{m=0}^{\infty} m \frac{\lambda^{m-1}}{(m-1)!} e^{-\lambda} - \lambda^2 = \\ &= e^{-\lambda} \cdot \lambda^2 \sum_{m=0}^{\infty} \frac{\lambda^{m-2}}{(m-2)!} + \lambda \sum_{m=0}^{\infty} e^{-\lambda} \frac{\lambda^{m-1}}{(m-1)!} - \lambda^2 = \lambda. \end{aligned} \quad (5.7)$$

Розподіл Пуассона залежить від одного параметра λ , який є одночасно математичним сподіванням і дисперсією випадкової величини X , розподіленої за законом Пуассона:

$$\mu = \sigma^2 = \lambda. \quad (5.8)$$

На рис. 5.3 наведений розподіл Пуассона при різних значеннях λ .

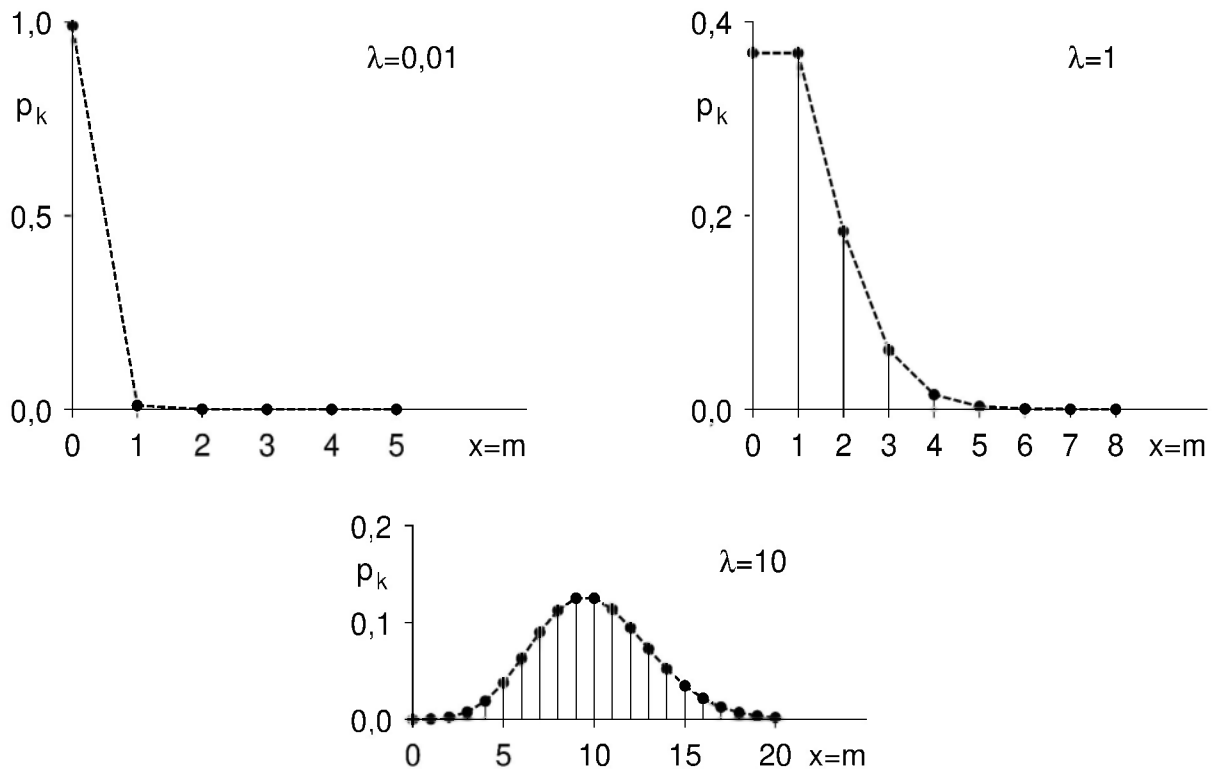


Рис. 5.3. Графіки розподілу Пуассона залежно від параметра λ

На практиці кількість дослідів n завжди обмежена, і важко перевірити виконання умови $\lambda_n = np_n \rightarrow \lambda$ при $n \rightarrow \infty$. Однак із цієї умови випливає, що $p_n = \lambda_n/n \rightarrow 0$. Тому якщо в схемі Бернуллі кількість дослідів n обмежена, але дуже велика, а ймовірність p дуже мала, тобто в кожному окремому досліді подія A настає дуже рідко, то як параметр беруть $\lambda=np$ і заміняють біномний розподіл $b(m;n,p)$ розподілом Пуассона $P_\lambda(m)$. Тому закон Пуассона ще називають **законом рідкісних подій**.

Так, наприклад, багато хвороб, на щастя, досить рідкісні або стають такими після здійснення профілактичних і лікувальних заходів. Однак навіть за найбільш сприятливих умов у великих популяціях усе-таки зустрічається деяка кількість хворих рідкісними хворобами. Якщо в популяції, що обстежується, наприклад, у конкретному місті, хворі зустрічаються частіше, ніж це прогнозується законом Пуассона для всієї країни, то це свідчить про порушення нормальних умов у даному місті, про необхідність з'ясування причин цього та вживання термінових заходів і т.д.

Приклад 5.3. (Гільдерман, 1991). При введенні вакцини проти поліомієліту імунітет створюється в 99,99 % випадків. Яка ймовірність того, що з 10 000 вакцинованих дітей:

- 1) не занедужає жодна дитина;
- 2) занедужає одна, дві, три, чотири, п'ять дітей?

Ймовірність занедужати $p = 0,0001$, кількість випробувань $n = 10\,000$. Тому $\lambda = np = 1$, за формулою Пуассона

$$P_{\lambda=1}(m=0) = \frac{1^0}{0!} e^{-1} \approx 0,368; \quad P_{\lambda=1}(m=3) = \frac{1^3}{3!} e^{-1} \approx 0,061;$$

$$P_{\lambda=1}(m=1) = \frac{1^1}{1!} e^{-1} \approx 0,368; \quad P_{\lambda=1}(m=4) = \frac{1^4}{4!} e^{-1} \approx 0,015;$$

$$P_{\lambda=1}(m=2) = \frac{1^2}{2!} e^{-1} \approx 0,184 \quad P_{\lambda=1}(m=5) = \frac{1^5}{5!} e^{-1} \approx 0,001$$

Якщо вважати, що народження 10 000 немовлят – це річна норма великого районного пологового будинку, то приблизно в 73% таких пологових будинків поліомієлітом занедужає не більше однієї дитини на рік ($P_{\lambda=1}(m=0) + P_{\lambda=1}(m=1) \approx 2 \times 0,368 = 0,736$); у 18% – дві дитини; у 92% – не більше двох дітей ($P_{\lambda=1}(m=0) + P_{\lambda=1}(m=1) + P_{\lambda=1}(m=2) \approx 0,368 + 0,368 + 0,184 = 0,92$); у 6% – три й в 1,5% – чотири дитини за рік. Якщо ж у деякому пологовому будинку занедужає п'ять дітей, це буде надзвичайна подія ($P_{\lambda=1}(m=5) \approx 0,001$). ●

5.3. Розподіл Бенфорда

Розподіл дискретної випадкової величини, відомий також як “закон аномальних чисел” або “закон розсіювання”, був відкритий американським фізиком Ф. Бенфордом у 1938 р. Початковий поштовх його пошукам дали бібліотечні таблиці логарифмів. Бенфорд помітив, що перші кілька сторінок більше забруднені, ніж наступні, тоді як останні майже зовсім чисті. Бенфорд зацікавився, чому інженери й студенти частіше користуються першими сторінками таблиць логарифмів, ніж останніми. Він вирішив, що це відбувається тому, що частіше доводиться мати справу з числами, які починаються з цифри 1, ніж із тими, що починаються з цифр 2, 3, 4 і т.д. Цими числами можуть бути й ті, які розміщені в спеціальних довідниках. Бенфорд дослідив близько 20 таблиць, у яких були дані про площу поверхні 335 рік, питому теплоємність і молярну масу тисяч хімічних сполук і навіть номери будинків перших 342 осіб, зазначених у біографічному довіднику американських учених. У результаті аналізу понад 20 тисяч чисел, що містилися в таблицях, була встановлена дивна закономірність. На перший погляд, здається правдоподібним припущення про рівноможливість для кожної з дев'яти цифр бути на першому місці в записі деякого числа. Отже, імовірність того, що число починається з деякої цифри, дорівнює $1/9 \approx 0,111^3$. Насправді ймовірність того, що число починається з цифри 1, виявилася майже в три рази більшою. Навпаки, імовірність того, що число починається з цифри 9, – майже в 2,5 рази меншою. Можна сказати, що в таблицях логарифмів, на які звернув увагу Бенфорд, перші сторінки виявилися брудніші останніх майже в сім разів ($3 \times 2,5 = 7,5$).

³ Виконуються всі умови для класичної ймовірності – є повна група несумісних і рівноможливих подій.

Бенфорд виразив знайдену ним закономірність математично у вигляді **закону аномальних чисел**: імовірність того, що випадковий десятковий дріб починається з цифри k ($k = \overline{1, K}$, $K = 9$), дорівнює

$$P_k = \lg(k+1) - \lg k = \lg \frac{k+1}{k}. \quad (5.9)$$

Для порівняння в табл. 5.7 і на рис. 5.4 наведені дослідні (вибіркові) дані Бенфорда p_k^* й теоретичні p_k .

Таблиця 5.7

K	1	2	3	4	5	6	7	8	9
p_k^*	0,306	0,185	0,124	0,094	0,080	0,064	0,051	0,049	0,047
p_k	0,301	0,176	0,125	0,097	0,079	0,067	0,058	0,051	0,046

Слово “аномальність” у назві закону з’явилося тому, що одні групи даних краще узгоджені з формулою (5.9), ніж інші. Площі рік і випадкові номери будинків дають кращий збіг із формулою (5.9), а таблиці квадратних коренів або питомих теплоємностей – гірший. Тому Бенфорд вирішив, що відкритий ним закон застосовується тільки до таких (аномальних) чисел, між якими не вбачається ніяких спільних для них закономірностей.

Пояснення закону Бенфорда таке. У певній множині чисел, що характеризують яке-небудь явище навколишнього світу, є найбільше число. Наприклад, кількість будинків на одній вулиці не перевищує загальної кількості будинків у даному населеному пункті, а площа поверхні будь-якої ріки – площі поверхні острова або материка, по якому вона протікає. Якщо в даному населеному пункті є кілька вулиць, то номери будинків на найбільш короткій із них будуть зустрічатися в ньому частіше інших номерів. У свою чергу, великих рік менше, ніж малих. Тому суть закону Бенфорда полягає в тому, що внаслідок обмеженості будь-яких множин числа, що їх характеризують, зустрічаються тим частіше, чим вони менші.

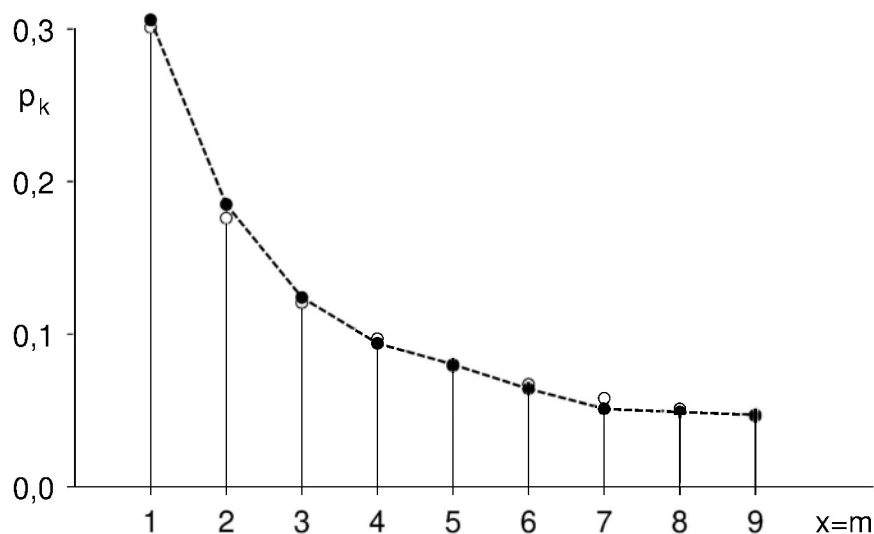


Рис. 5.4. Розподіл Бенфорда (білі кружечки – теоретичний, чорні – емпіричний)

6. Стандартні розподіли неперервних випадкових величин

6.1. Рівномірний розподіл

Неперервна випадкова величина X має рівномірний розподіл на відрізку $[a, b]$, якщо на цьому відрізку її функція густини є стала, а поза ним дорівнює нулю (рис. 6.1):

$$f(x) = \begin{cases} 0, & x < a, \\ c = \text{const}, & a \leq x \leq b, \\ 0, & x > b. \end{cases} \quad (6.1)$$

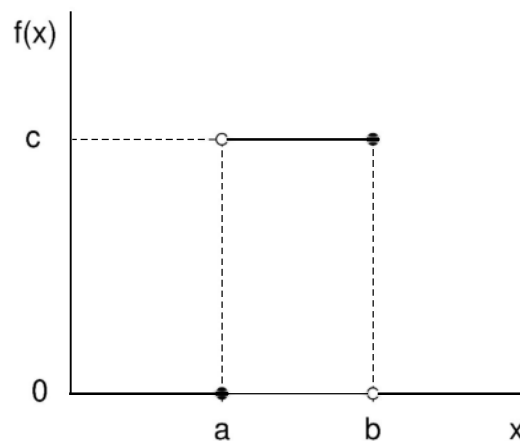


Рис. 6.1. Графік функції густини рівномірного розподілу

Оскільки площа, обмежена кривою густини, дорівнює одиниці (див. властивість 2 в розд. 4.8), то

$$\int_{-\infty}^{+\infty} f(x) dx = \int_{-\infty}^a f(x) dx + \int_a^b f(x) dx + \int_b^{+\infty} f(x) dx = \int_a^b f(x) dx = \int_a^b c dx = 1,$$

звідки $c = 1/(b-a)$. Отже, остаточно закон рівномірного розподілу можна записати у вигляді

$$f(x) = \begin{cases} 0, & x < a, \\ 1/(b-a), & a \leq x \leq b, \\ 0, & x > b. \end{cases} \quad (6.2)$$

Застосовуючи формулу (6.2) і зв'язок між функцією густини $f(x)$ і функцією розподілу $F(x)$, знайдемо вираз функції розподілу $F(x)$ для рівномірного розподілу:

$$F(x) = \int_{-\infty}^x f(t) dt = \begin{cases} 0, & x < a, \\ \int_a^x f(t) dt, & a \leq x \leq b, \\ 1, & x > b. \end{cases}$$

Оскільки $\int_a^x f(t) dt = \int_a^x \frac{dt}{b-a} = \frac{x-a}{b-a}$, то остаточно

$$F(x) = \begin{cases} 0, & x < a, \\ \frac{x-a}{b-a}, & a \leq x \leq b, \\ 1, & x > b. \end{cases} \quad (6.3)$$

Графік функції $F(x)$ показаний на рис. 6.2.

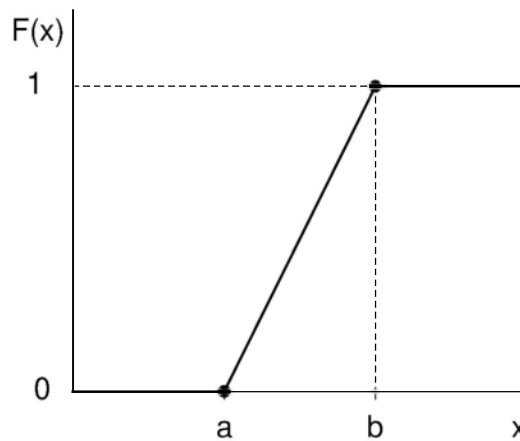


Рис. 6.2. Графік функції рівномірного розподілу

Визначимо математичне сподівання й дисперсію випадкової величини X , що має рівномірний розподіл:

$$\mu = M[X] = \int_a^b x f(x) dx = \int_a^b \frac{x dx}{b-a} = \frac{b^2 - a^2}{2(b-a)} = \frac{a+b}{2}, \quad (6.4)$$

Дисперсія та середньоквадратичне відхилення X відповідно дорівнюють:

$$\sigma^2 = D[X] = \int_a^b (x - \mu)^2 f(x) dx = \int_a^b \left(x - \frac{a+b}{2}\right)^2 \frac{dx}{b-a} = \frac{(b-a)^2}{12}, \quad (6.5)$$

$$\sigma = (b - a) / 2\sqrt{3}. \quad (6.6)$$

Ймовірність належності значень випадкової величини X , що має рівномірний розподіл, проміжку $[\alpha, \beta] \in [a, b]$

$$P(\{\alpha < x < \beta\}) = \int_{\alpha}^{\beta} \frac{dx}{b-a} = \frac{\beta - \alpha}{b-a}. \quad (6.7)$$

6.2. Показниковий розподіл

Неперервна випадкова величина X розподілена за показниковим (експонентним) законом із параметром α , якщо її функція густини розподілу є така:

$$f(x) = \begin{cases} 0, & x < 0, \\ \alpha \cdot e^{-\alpha x}, & x \geq 0. \end{cases} \quad (6.8)$$

Функцію розподілу $F(x)$ одержуємо безпосереднім інтегруванням функції $f(x)$:

$$F(x) = \int_{-\infty}^x f(x) dx = \begin{cases} \int_0^x \alpha e^{-\alpha} dx = -e^{-\alpha} x \Big|_0^x = 1 - e^{-\alpha x}, & x > 0, \\ 0, & x \leq 0. \end{cases}$$

Графіки функції густини і розподілу показані на рис. 6.3 і 6.4.

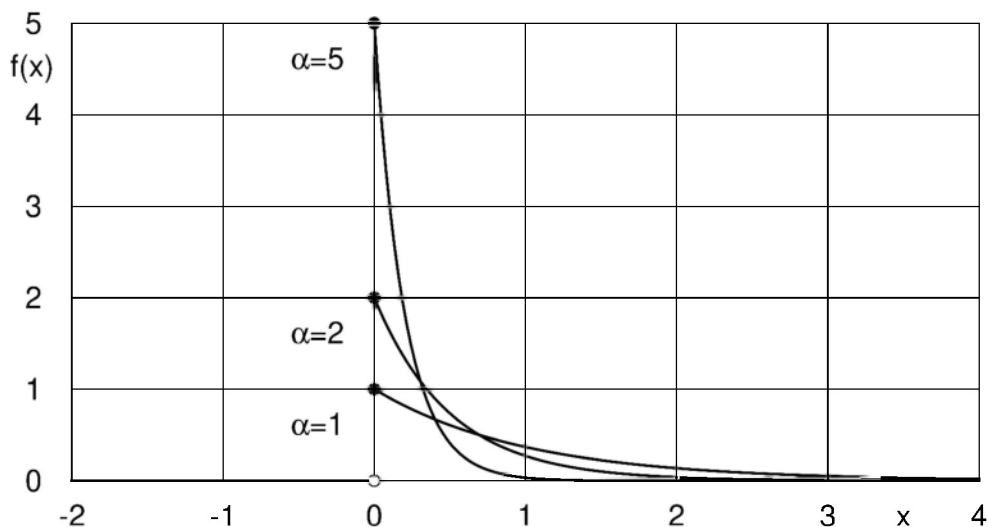


Рис. 6.3. Графік функції густини показникового розподілу залежно від значення α

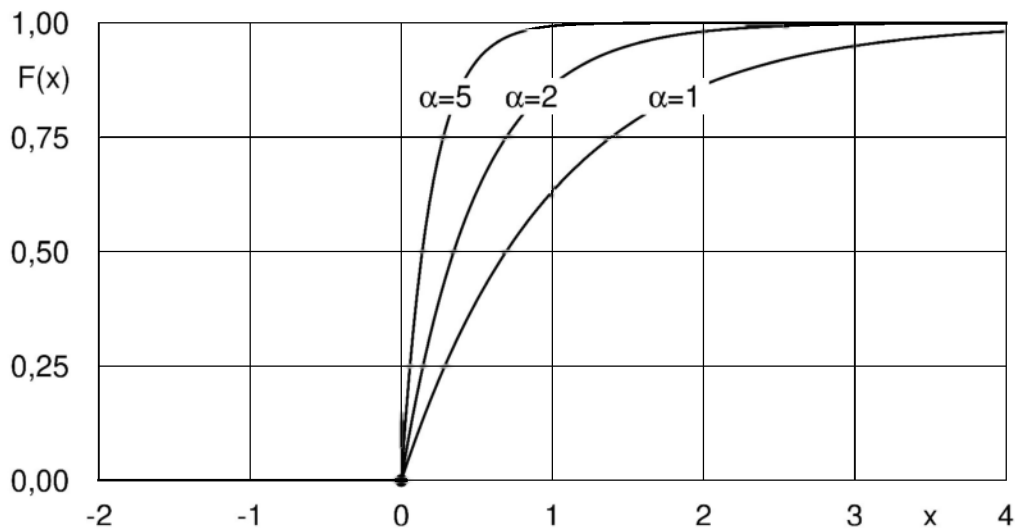


Рис. 6.4. Графік функції густини показникового розподілу залежно від значення α

Математичне сподівання показникового розподілу дорівнює

$$\mu = M[X] = \int_0^{+\infty} x \cdot \alpha e^{-\alpha x} dx = \int_0^{+\infty} e^{-\alpha x} dx = -\frac{1}{\alpha} e^{-\alpha x} \Big|_0^{+\infty} = \frac{1}{\alpha}, \quad (6.9)$$

а дисперсія –

$$\sigma^2 = D[X] = \int_0^{+\infty} x^2 \cdot \alpha e^{-\alpha x} dx - \frac{1}{\alpha^2} = \frac{2}{\alpha^2} - \frac{1}{\alpha^2} = \frac{1}{\alpha^2}. \quad (6.10)$$

6.3. Нормальний розподіл

Неперервна випадкова величина X розподілена за нормальним законом із параметрами a , b , якщо її функція густини є така:

$$f(x) = \frac{1}{b\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{x-a}{b}\right)^2}, \quad (6.11)$$

$$-\infty < x < +\infty, \quad -\infty < a < +\infty, \quad b > 0.$$

Подивимося, як буде змінюватися графік функції густини розподілу залежно від зміни значень параметрів a і b . При зміні a крива $f(x)$, не змінюючи своєї форми, буде зміщуватись уздовж осі абсцис (рис. 6.5).

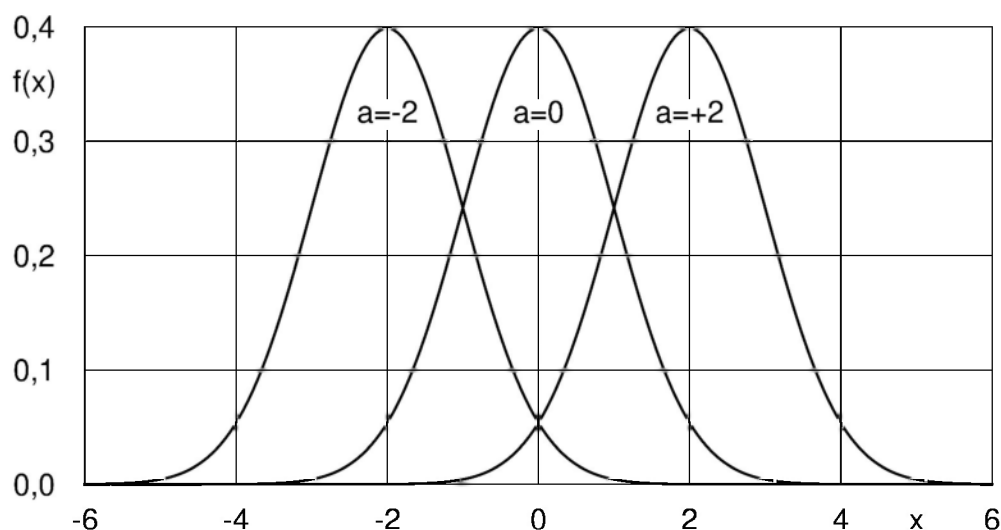


Рис. 6.5 Графік функції густини нормального розподілу залежно від значення a

Зміна b рівносильна зміні масштабу кривої за обома осями. Наприклад, при подвоєнні b масштаб по осі абсцис подвоїться, а по осі ординат – зменшиться у два рази (рис. 6.6). Крива густини при цьому стає більш плоскою, розтягуючись уздовж осі абсцис. При зменшенні b крива витягається вгору, одночасно стискаючись із боків. В обох випадках крива обмежує зверху осі у одиничну площу (властивість 2 функції розподілу, див. розд. 4.8):

$$\int_{-\infty}^{+\infty} \frac{1}{b\sqrt{2\pi}} \exp\left\{-\frac{1}{2}\left(\frac{x-a}{b}\right)^2\right\} dx = 1.$$

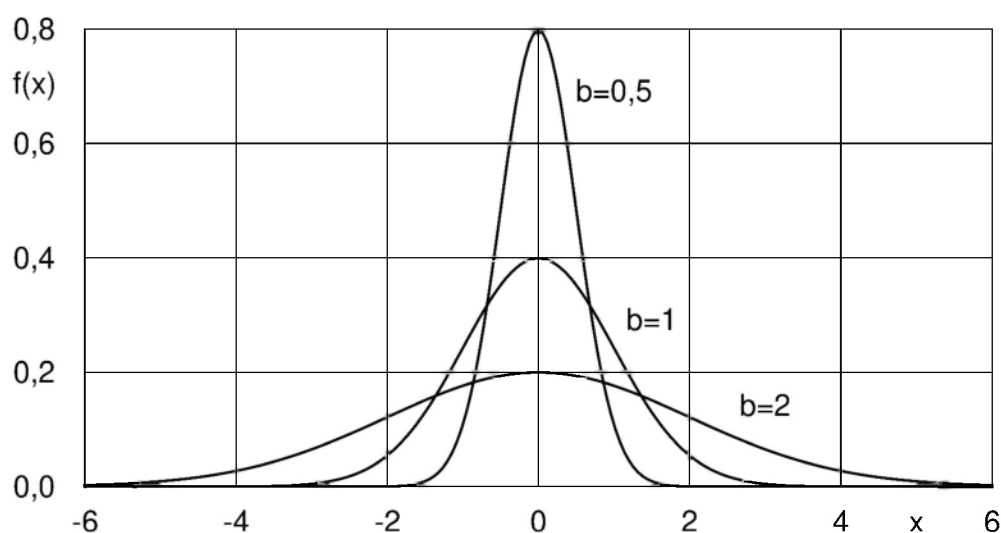


Рис. 6.6 Графік функції густини показникового розподілу залежно від значення b

Обчислимо основні характеристики випадкової величини Y , розподіленої за нормальним законом. Почнемо з математичного сподівання:

$$\mu = \int_{-\infty}^{+\infty} x f(x) dx = \frac{1}{b\sqrt{2\pi}} \int_{-\infty}^{+\infty} x \exp\left(-\frac{1}{2}\left(\frac{x-a}{b}\right)^2\right) dx.$$

Застосовуючи заміну змінної $y = \frac{x-a}{b\sqrt{2}}$, одержимо

$$\mu = \frac{1}{\sqrt{\pi}} \int_{-\infty}^{+\infty} (b\sqrt{2} + a) e^{-y^2} dy = \frac{b\sqrt{2}}{\pi} \int_{-\infty}^{+\infty} y e^{-y^2} dy + \frac{a}{\sqrt{\pi}} \int_{-\infty}^{+\infty} e^{-y^2} dy.$$

Перший з інтегралів дорівнює нулю як інтеграл у симетричних межах від непарної функції; другий являє собою відомий **інтеграл Ейлера – Пуассона**:

$$\int_{-\infty}^{+\infty} e^{-y^2} dy = 2 \int_0^{+\infty} e^{-y^2} dy = \sqrt{\pi}. \quad (6.12)$$

Отже математичне сподівання величини Y дорівнює a :

$$\mu = M[x] = a, \quad (6.13)$$

тобто параметр a у виразі (6.11) є математичним сподіванням.

Обчислимо дисперсію випадкової величини X :

$$\sigma^2 = \frac{1}{b\sqrt{2\pi}} \int_{-\infty}^{+\infty} (x-a)^2 \exp\left(-\frac{1}{2}\left(\frac{x-a}{b}\right)^2\right) dx.$$

Застосуємо заміну змінної $x \rightarrow y$:

$$\sigma^2 = \frac{2b^2}{\sqrt{\pi}} \int_{-\infty}^{+\infty} y^2 e^{-y^2} dy.$$

Інтегруючи отриманий вираз частинами, одержимо

$$\sigma^2 = \frac{b^2}{\sqrt{\pi}} \int_{-\infty}^{+\infty} 2y^2 e^{-y^2} dy = \frac{b^2}{\sqrt{\pi}} \left\{ -y e^{y^2} \Big|_{-\infty}^{+\infty} + \int_{-\infty}^{+\infty} e^{y^2} dy \right\}.$$

Перший доданок у фігурних дужках дорівнює нулю, тому що $\exp(-y^2)$ при $y \rightarrow \infty$ убуває швидше, ніж зростає будь-який степінь y , другий доданок дорівнює $\sqrt{\pi}$, звідки

$$\sigma^2 = D[x] = b^2. \quad (6.14)$$

Отже, параметр b у виразі (6.12) є не що інше, як середнє квадратичне відхилення σ випадкової величини X . Остаточно функція густини нормального розподілу записується в такому вигляді:

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2}. \quad (6.15)$$

Якщо неперервна випадкова величина розподілена за нормальним законом (6.12), (6.13) із відомими параметрами $a = \mu$, $b = \sigma$, можна використовувати скорочений запис $X \in N(\mu, \sigma)$. Якщо внаслідок нормалізації одержана випадкова величина $Z = (x - \mu)/\sigma$, то легко впевнитися, що для неї $\mu_z = 1$, $\sigma_z = 0$, що скорочено записується в такий спосіб: $Z \in N(0, 1)$.

Обчислимо для нормально розподіленої випадкової величини ймовірність набуття значень у проміжку $[x_1, x_2]$:

$$P(\{x_1 < X < x_2\}) = \int_{x_1}^{x_2} f(x) dx = \frac{1}{\sigma\sqrt{2\pi}} \int_{x_1}^{x_2} \exp\left(-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2\right) dx.$$

Замінивши в інтегралі змінну $z = (x - \mu)/\sigma$ (така заміна називається нормалізацією) і відповідно змінивши межі інтегрування, одержимо

$$P(\{x_1 < X < x_2\}) = \frac{1}{\sqrt{2\pi}} \int_{\frac{x_1-\mu}{\sigma}}^{\frac{x_2-\mu}{\sigma}} e^{-z^2/2} dz.$$

Невизначений інтеграл $\int e^{-z^2/2} dz$ не виражається через елементарні функції, але його можна виразити через спеціальну функцію

$$\Phi(z) = \frac{1}{\sqrt{2\pi}} \int_0^z e^{-z^2/2} dz, \quad (6.16)$$

названу **функцією Лапласа** (або інтегралом ймовірностей). Значення цієї функції наведені в додатку. За допомогою цієї функції шукана ймовірність виражається простою формулою:

$$P(\{x_1 < X < x_2\}) = \Phi\left(\frac{y_2 - \mu}{\sigma}\right) - \Phi\left(\frac{y_1 - \mu}{\sigma}\right).$$

Властивості функції Лапласа

1. $\Phi(0) = 0$.

2. $\Phi(-z) = -\Phi(z)$ (непарна функція).

Для доведення достатньо зробити заміну змінної $t = -z$:

$$\Phi(-z) = \frac{1}{\sqrt{2\pi}} \int_0^{-z} e^{-x^2/2} dx = -\frac{1}{\sqrt{2\pi}} \int_0^z e^{-t^2/2} dt = -\Phi(z). \blacksquare \quad (6.17)$$

3. $\Phi(\pm \infty) = \pm 0,5$.

Властивість перевіряється безпосередньо застосуванням інтеграла Ейлера – Пуассона (6.12):

$$\frac{1}{\sqrt{2\pi}} \int_0^{+\infty} e^{-z^2/2} dz = \frac{1}{\sqrt{\pi}} \int_0^{+\infty} e^{-(z/\sqrt{2})^2} d(z/\sqrt{2}) = \frac{1}{\sqrt{\pi}} \int_0^{+\infty} e^{-y^2} dy = \frac{1}{\sqrt{\pi}} \frac{\sqrt{\pi}}{2} = \frac{1}{2}. \blacksquare \quad (6.18)$$

Найбільш просто через функцію Лапласа (6.18) виражається ймовірність потрапляння нормально розподіленої випадкової величини Y на відрізок довжиною $2h$, симетричний щодо математичного сподівання μ :

$$P(\{\mu - h < Y < \mu + h\}) = P(\{|Y - \mu| < h\}) = \Phi\left(\frac{\mu - \mu + h}{\sigma}\right) - \Phi\left(\frac{\mu - \mu - h}{\sigma}\right),$$

або з урахуванням властивості 2:

$$P(\{|Y - \mu| < h\}) = 2 \Phi\left(\frac{h}{\sigma}\right). \quad (6.19)$$

Через функцію Лапласа виражається й функція розподілу $F(y)$ випадкової величини Y . Вважаючи, що $x_1 \rightarrow -\infty$, $x_2 = x$ і з огляду на те, що $\Phi(-\infty) = -0,5$, одержимо таке:

$$F(x) = \frac{1}{2} + \Phi\left(\frac{x - \mu}{\sigma}\right) = \frac{1}{2} + \Phi(z).$$

Отже, за допомогою функції розподілу $F(x)$ (6.19) або функції Лапласа $\Phi(z)$ (6.16) можна визначити ймовірність набуття випадковою величиною X , розподіленою за нормальним законом (6.11), (6.15) із відомими параметрами $a = \mu$, $b = \sigma$ значень із певного проміжку. Якщо звернутися до механічної моделі неперервної випадкової величини, то цю ймовірність можна розглядати як частку об'єктів генеральної сукупності, у яких значення числової ознаки X належить цьому проміжку. Для проміжків $\mu \pm 1\sigma$, $\mu \pm 2\sigma$, $\mu \pm 3\sigma$ ці частки дорівнюють відповідно:

$$\begin{aligned} P(\mu - 1\sigma \leq z \leq \mu + 1\sigma) &= 2 \Phi(1) = 0,6827; \\ P(\mu - 2\sigma \leq z \leq \mu + 2\sigma) &= 2 \Phi(2) = 0,9545; \\ P(\mu - 3\sigma \leq z \leq \mu + 3\sigma) &= 2 \Phi(3) = 0,9973, \end{aligned}$$

що становить зміст правила 3σ (правила “трьох сигм”). Графічна ілюстрація правила 3σ дана на рис. 6.7.

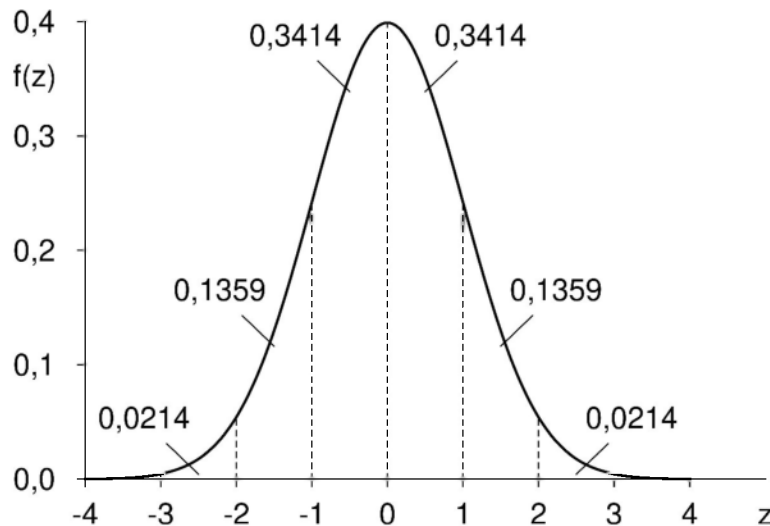


Рис. 6.7. Графічне пояснення правила 3σ

Нормальний розподіл є одним з основних у теорії випадкових величин і статистиці, оскільки він є гарним наближенням до багатьох фізіологічних показників (зріст, маса тіла, обхват грудей, стегон, пульс, артеріальний і венозний тиск).

За допомогою правила 3σ для об'єктів, що є носіями нормально розподіленої ознаки, звичайно вводиться класифікація, наведена в табл.6.1 (межі класифікації дані для нормалізованої ознаки $Z = (X - \mu)/\sigma$).

Таблиця. 6.1

z	Значення ознаки
$z < -2$	Дуже мале
$-2 \leq z < -1$	Мале
$-1 \leq z < +1$	Нормальне
$+1 \leq z < +2$	Велике
$2 < z$	Дуже велике

Приклад 6.1. (Плохинський, 1962). Для хлопчика 14 років необхідно зробити оцінку статури за даними для цієї вікової групи (табл. 6.2).

Таблиця 6.2

Показник	X	μ	σ	Z
Зріст, см	159	142	8,5	+2,0
Обхват грудей, см	64	66	4,0	- 0,5
Обхват голови, см	52	50	2,0	+1,0
Маса тіла, кг	40	40	6,0	0,0
Ємність легенів, см ³	2 070	2 300	410	- 0,5

Очевидно, що це високий вузькогрудий великоголовий хлопчик астеничного типу. ●

6.4. Функція Лапласа

Розглянемо застосування прийомів роботи з функцією Лапласа на прикладі розв'язання прямої й оберненої задач для квантилів неперервної випадкової величини $X \sim N(\mu, \sigma)$.

Алгоритм розв'язання прямої задачі

У прямій задачі для двох значень x_1 і x_2 неперервної випадкової величини X потрібно обчислити ймовірність того, що вона набуває значень між x_1 і x_2 , тобто $P(\{x_1 \leq X \leq x_2\})$. Шукана величина може бути інтерпретована як частка об'єктів генеральної сукупності, у яких значення ознаки X знаходиться між x_1 і x_2 . При відомому обсязі генеральної сукупності N може бути обчислена й кількість таких об'єктів як добуток N і шуканої ймовірності.

Оскільки функція Лапласа введена на основі стандартизованого нормального розподілу, то стандартизуються й значення x_1 і x_2 :

$$z_1 = \frac{x_1 - \mu}{\sigma}, \quad z_2 = \frac{x_2 - \mu}{\sigma}. \quad (6.20)$$

Розглянемо можливі ситуації в прямій задачі, обумовлені вибором нормалізованих значень z_1 і z_2 :

а) $0 \leq z_1 \leq z_2$ (рис. 6.8).

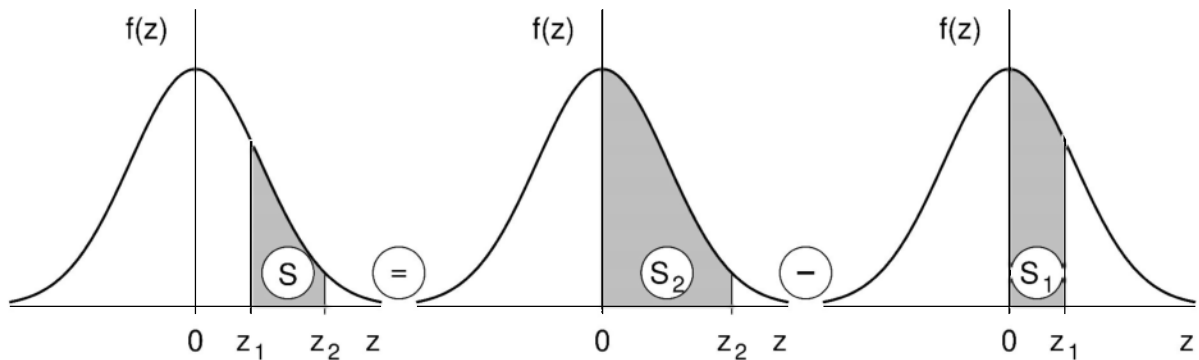


Рис. 6.8. Пряма задача для додатних значень z_1 і z_2

Графічний розв'язок, наведений на рис.6.8, показує, що $P(\{x_1 \leq X \leq x_2\}) = S = S_2 - S_1$. Шукана ймовірність дорівнює площі S під кривою $f(z)$, яка обмежена вертикальними прямими $z = z_1$ і $z = z_2$, і може бути обчислена як різниця площі $S_2 = \Phi(z_2)$ під кривою $f(z)$ між прямими $z = 0$ і $z = z_2$ і площі $S_1 = \Phi(z_1)$ під кривою $f(z)$ між прямими $z = 0$ і $z = z_1$.

Точний розв'язок задачі в цьому випадку записується так:

$$P(\{x_1 \leq X \leq x_2\}) = P(\{z_1 \leq \tilde{z} \leq z_2\}) = \Phi(z_2) - \Phi(z_1). \quad (6.21)$$

У додатку потрібно знайти найближчі до z_1 і z_2 (6.20) значення, нехай це будуть z_a і z_b . Тоді замість точного розв'язку (6.21) можна записати наближений:

$$P(\{x_1 \leq X \leq x_2\}) = P(\{z_1 \leq \tilde{z} \leq z_2\}) \approx \Phi(z_b) - \Phi(z_a); \quad (6.22)$$

б) $z_1 \leq 0 \leq z_2$ (рис. 6.9).

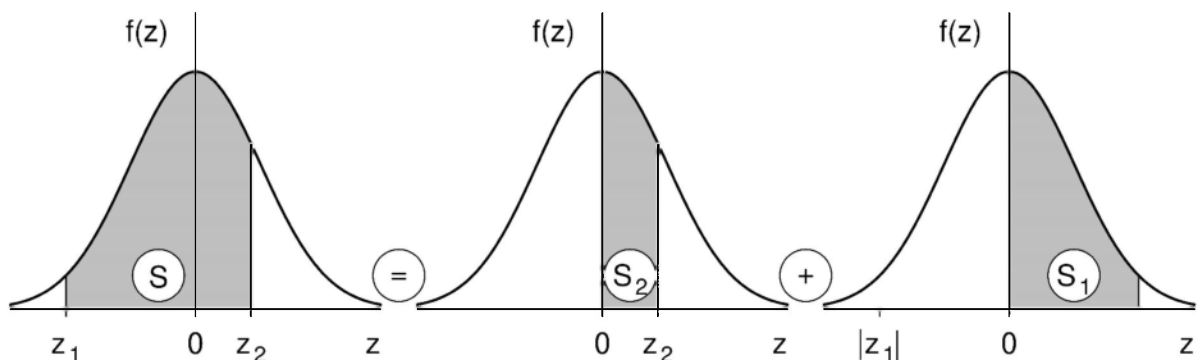


Рис. 6.9. Пряма задача для від'ємного значення z_1 і додатного значення z_2

Графічне розв'язання зводиться до знаходження площі криволінійної трапеції $S = S_1 + S_2$. Однак розв'язок (6.21) тут не може бути застосований, тому що в додатку наведені значення $\Phi(z)$ тільки для додатних значень z , отже, $\Phi(z_1)$ обчислити не можна. У такому випадку варто застосувати властивість 2 функції Лапласа, із якої випливає, що $S_1 = \Phi(|z_1|)$. Оскільки $z_2 > 0$, то $S_2 = \Phi(z_2)$. Тому точний розв'язок буде записаний так:

$$P(\{x_1 \leq X \leq x_2\}) = P(\{z_1 \leq \tilde{z} \leq z_2\}) = \Phi(|z_1|) - \Phi(|z_2|). \quad (6.23)$$

Як і в попередньому випадку, потрібно знайти в додатку найближчі до $|z_1|$ і z_2 значення. Нехай це будуть z_a і z_b відповідно. Тоді наближений розв'язок прямої задачі запишеться так:

$$P(\{x_1 \leq X \leq x_2\}) = P(\{z_1 \leq \tilde{z} \leq z_2\}) \approx \Phi(z_a) - \Phi(z_b); \quad (6.24)$$

в) $z_1 < z_2 \leq 0$ (рис. 6.10).

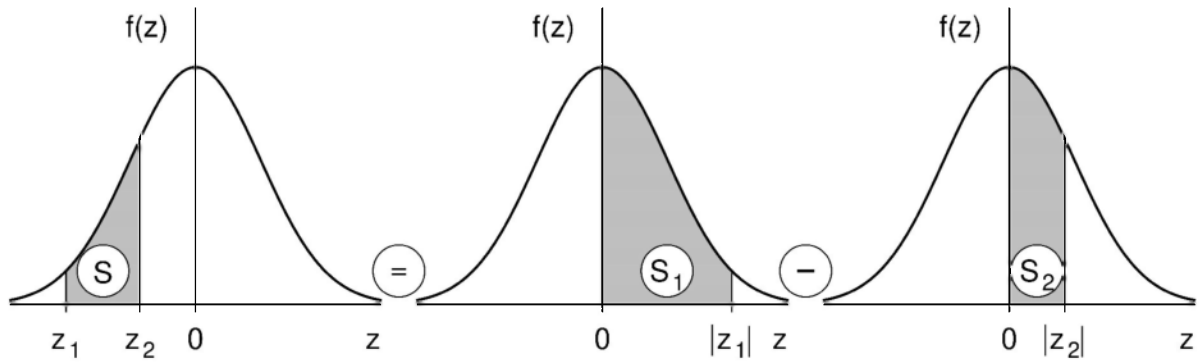


Рис. 6.10. Пряма задача для від'ємних значень z_1 і z_2

Графічне розв'язання зводиться до знаходження площі $S = S_1 - S_2$. Як і в попередньому випадку, потрібно враховувати, що значення z_1 і z_2 відсутні в додатку (крім випадку $z_2 = 0$). Тому знову варто застосувати властивість 2, із якої одержуємо, що $S_1 = \Phi(|z_1|)$, $S_2 = \Phi(|z_2|)$. Точний розв'язок у цьому випадку записується в такому вигляді:

$$P(\{x_1 \leq X \leq x_2\}) = P(\{z_1 \leq \tilde{z} \leq z_2\}) = \Phi(|z_1|) - \Phi(|z_2|). \quad (6.25)$$

Для запису наближеного розв'язку, що відповідає (6.25), потрібно знайти в додатку значення z_a і z_b , найближчі до $|z_1|$ і $|z_2|$:

$$P(\{x_1 \leq X \leq x_2\}) = P(\{z_1 \leq \tilde{z} \leq z_2\}) \approx \Phi(|z_a|) - \Phi(|z_b|); \quad (6.26)$$

г) $0 \leq z_1 \leq \infty$ (рис. 6.11).

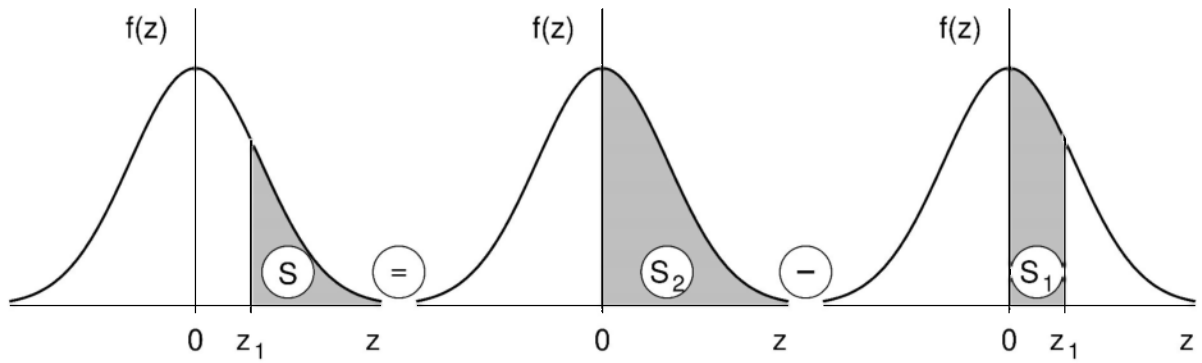


Рис. 6.11. Пряма задача для нескінченного додатного значення z_2

Це розглянута вище ситуація (а) при $z_2 \rightarrow \infty$. Оскільки $\Phi(z_2) = \Phi(+\infty) = 0,5$, точний розв'язок отримуємо з формули (6.22):

$$P(\{x_1 \leq X \leq x_2\}) = P(\{z_1 \leq \tilde{z} \leq z_2\}) = \frac{1}{2} - \Phi(z_1), \quad (6.27)$$

а наближений – із виразу (6.22):

$$P(\{x_1 \leq X \leq x_2\}) = P(\{z_1 \leq \tilde{z} \leq z_2\}) \approx \frac{1}{2} - \Phi(z_a); \quad (6.28)$$

д) $-\infty < z_2 \leq 0$ (рис. 6.12).

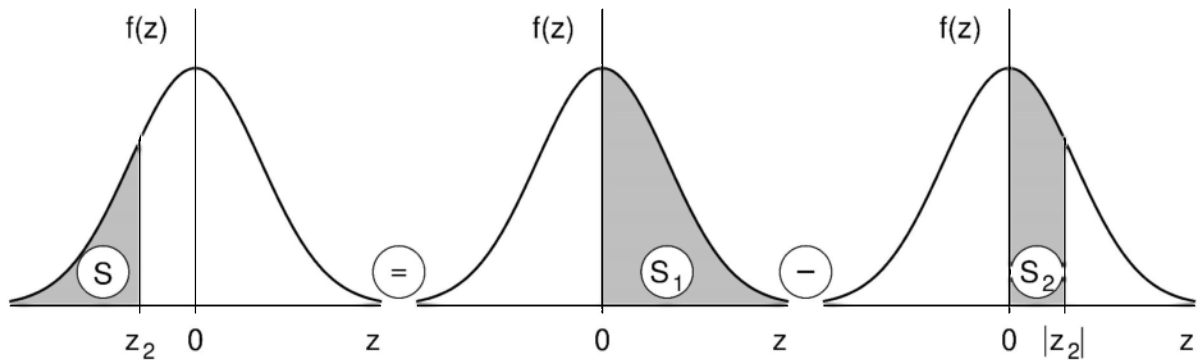


Рис. 6.12. Пряма задача для нескінченного від'ємного значення z_1

Це розглянута вище ситуація (в) при $z_1 \rightarrow -\infty$. Оскільки $\Phi(|z_1|) = \Phi(+\infty) = 0,5$, точний розв'язок отримаємо з виразу (6.25):

$$P(\{x_1 \leq X \leq x_2\}) = P(\{z_1 \leq \tilde{z} \leq z_2\}) = \frac{1}{2} - \Phi(|z_2|), \quad (6.29)$$

а наближений – із формули (6.26):

$$P(\{x_1 \leq X \leq x_2\}) = P(\{z_1 \leq \tilde{z} \leq z_2\}) \approx \frac{1}{2} - \Phi(|z_b|); \quad (6.30)$$

е) $-\infty < z < +\infty$ (рис. 6.13).

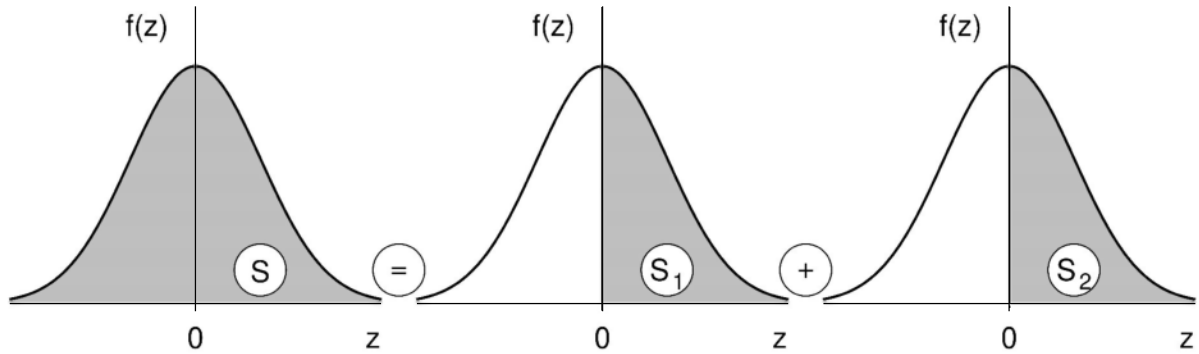


Рис. 6.13. Пряма задача для нескінченних значень z_1 і z_2

Цей випадок випливає з випадку (б) при $z_1 \rightarrow -\infty$, $z_2 \rightarrow +\infty$. Практично він не представляє інтересу, тому що шукана ймовірність, яка дорівнює одиниці, відома з властивості 2 функції густини, і наведений тут тільки з методичних міркувань. Дійсно, із формули (6.23) випливає, що

$$\begin{aligned} P(\{-\infty \leq X \leq +\infty\}) &= P(\{-\infty \leq \tilde{z} \leq +\infty\}) = \\ &= \lim_{|z_1 \rightarrow +\infty|} \Phi(|z_1|) + \lim_{|z_2 \rightarrow +\infty|} \Phi(|z_2|) = 0,5 + 0,5 = 1. \end{aligned} \quad (6.31)$$

Усі отримані розв'язки зведемо в табл. 6.3, причому помістимо тут і індекси цих розв'язків як квантилів.

Таблиця. 6.3.

z_1, z_2	q_1	q_2	$P(\{z_1 \leq \tilde{z} \leq z_2\}) = F(z_2) - F(z_1) = q_2 - q_1$
а) $0 \leq z_1 \leq z_2$	$0,5 + \Phi(z_1)$	$0,5 + \Phi(z_2)$	$\Phi(z_2) - \Phi(z_1)$
б) $z_1 \leq 0 \leq z_2$	$0,5 - \Phi(z_1)$	$0,5 + \Phi(z_2)$	$\Phi(z_1) + \Phi(z_2)$
в) $z_1 < z_2 \leq 0$	$0,5 - \Phi(z_1)$	$0,5 - \Phi(z_2)$	$\Phi(z_1) - \Phi(z_2)$
г) $0 \leq z_1 \leq \infty$	$0,5 + \Phi(z_1)$	1	$0,5 - \Phi(z_1)$
д) $-\infty < z_2 \leq 0$	0	$0,5 - \Phi(z_2)$	$0,5 - \Phi(z_2)$
е) $-\infty < z < +\infty$	0	1	1

Алгоритм розв'язання оберненої задачі

В оберненій задачі для заданої ймовірності $P(\{x_1 \leq X \leq x_2\}) > 0$ того, що неперервна випадкова величина $X \sim N(\mu, \sigma)$ набуває значень між x_1 і x_2 , потрібно знайти ці значення. Оскільки неперервна випадкова величина X , розподілена за нормальним законом, теоретично може набувати будь-яких значень на всій числовій прямій $-\infty < x_1 < x_2 < +\infty$, то звузимо постановку задачі, виходячи з практично спостережуваних значень безперервної випадкової величини X . Будемо вважати, що

$$P(\{x_1 \leq X \leq x_2\}) \leq P(\{\mu - 3\sigma \leq X \leq \mu + 3\sigma\}). \quad (6.32)$$

Із формули (6.32) випливає, що розв'язок оберненої задачі потрібно шукати в діапазоні значень $\mu \pm 3\sigma$, який фігурує в правилі 3σ (див. розд. 6.3).

Один розв'язок, що називається мінімальним, можна побудувати вправо від значення безперервної випадкової величини X , яка дорівнює $\mu - 3\sigma$. Для цього візьмемо $x_1 = \mu - 3\sigma$. Квантиль q_1 цього значення обчислюємо в такій послідовності:

1) виконуємо стандартизацію:

$$z_1 = \frac{x_1 - \mu}{\sigma} = \frac{\mu - 3\sigma - \mu}{\sigma} = -3;$$

2) за додатком $\Phi(|z_1|) = \Phi(3) \approx 0,49865$;

3) за табл. 6.3. розв'язків прямої задачі визначаємо індекс q_1 значення x_1 як квантиля, тобто

$$q_1 = 0,5 - \Phi(|z_1|) = 0,5 - \Phi(3) \approx 0,00135.$$

Очевидно, що індекс q_2 невідомого значення x_2 як квантиля дорівнює

$$q_2 = q_1 + P(\{x_1 \leq X \leq x_2\}).$$

Отже, значення x_2 знайдемо за його відомим квантилем q_2 . Для цього застосуємо таку послідовність дій:

1) якщо $q_2 \leq 0,5$, тоді $z_2 \leq 0$ (рис.6.14) і індекс q_2 згідно з табл. 6.3 зв'язаний із невідомим стандартизованим значенням z_2 співвідношенням

$$q_2 = 0,5 - \Phi(|z_2|),$$

звідки випливає, що потрібно знайти за додатком таке значення $|z_2|$, для якого

$$\Phi(|z_2|) = 0,5 - q_2,$$

після чого визначаємо $z_2 = -|z_2|$ і $x_2 = \mu + z_2\sigma$;

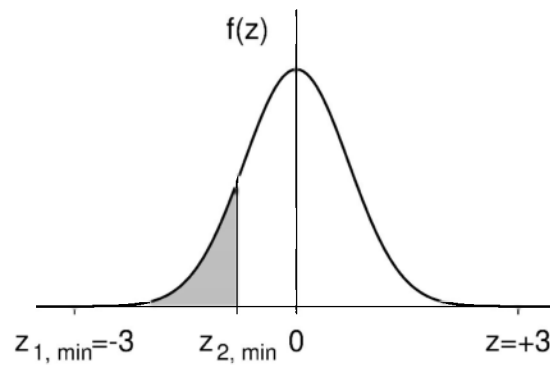


Рис. 6.14. Мінімальний розв'язок оберненої задачі для $q_2 \leq 0,5$, $z_2 \leq 0$

2) якщо $q_2 \geq 0,5$, тоді $z_2 \geq 0$ (рис. 6.15) і індекс q_2 згідно з табл. 6.3 зв'язаний із невідомим стандартизованим значенням z_2 співвідношенням

$$q_2 = 0,5 + \Phi(z_2),$$

звідки випливає, що потрібно знайти в додатку таке значення z_2 , для якого

$$\Phi(z_2) = q_2 - 0,5,$$

після чого знаходимо

$$x_2 = \mu + z_2\sigma.$$

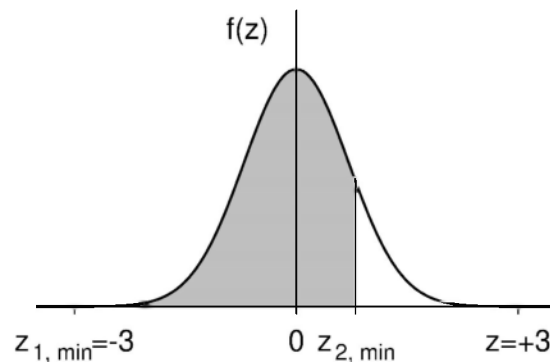


Рис. 6.15. Мінімальний розв'язок оберненої задачі для $q_2 \geq 0,5$, $z_2 \geq 0$

Отримані мінімальні розв'язки для початкової випадкової величини X та стандартизованої змінної Z позначимо відповідно як $(x_{1 \min}, x_{2 \min})$ та $(z_{1 \min}, z_{2 \min})$.

Інший розв'язок, що називається максимальним, можна побудувати вліво від значення неперервної випадкової величини X , яке дорівнює $\mu + 3\sigma$. Для цього візь-

мемо $x_2 = \mu + 3\sigma$. Індекс q_2 цього значення, що розглядається як квантиль, обчислюється в такій послідовності:

1) виконується стандартизація:

$$z_2 = \frac{x_2 - \mu}{\sigma} = \frac{\mu + 3\sigma - \mu}{\sigma} = +3;$$

2) за додатком $\Phi(z_2) = \Phi(3) \approx 0,49865$;

3) за табл.6.3 визначається індекс q_2 значення x_2 як квантиля, тобто

$$q_2 = 0,5 + \Phi(z_2) = 0,5 + \Phi(3) \approx 0,99865.$$

Очевидно, що індекс q_1 невідомого значення x_1 як квантиля дорівнює

$$q_1 = q_2 - P(\{x_1 \leq X \leq x_2\}).$$

Отже, значення x_1 слід знайти за його індексом q_1 . Для цього застосуємо таку послідовність дій:

1) якщо $q_1 \leq 0,5$, тоді $z_1 \leq 0$ (рис.6.16) і індекс згідно з табл.6.3 зв'язаний із невідомим стандартизованим значенням x_1 співвідношенням

$$q_1 = 0,5 - \Phi(|z_1|),$$

а отже, потрібно знайти в додатку таке значення $|z_1|$, для якого

$$\Phi(|z_1|) = 0,5 - q_1,$$

після чого буде знайдено $z_1 = -|z_1|$ і $x_1 = \mu + z_1\sigma$;

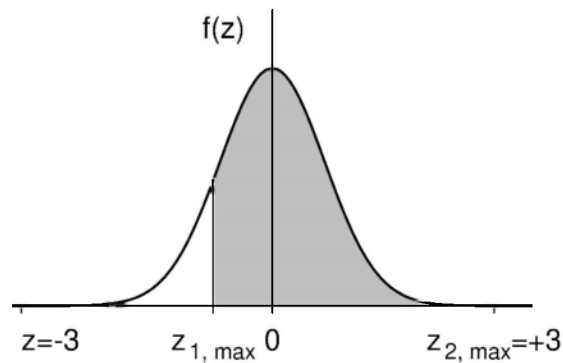


Рис. 6.16. Максимальний розв'язок оберненої задачі для $q_2 \leq 0,5$, $z_1 \leq 0$

2) якщо $q \geq 0,5$, тоді $z_1 \geq 0$ (рис 6.17) і індекс q_1 згідно з табл.6.3 зв'язаний із невідомим стандартизованим значенням z_1 співвідношенням

$$q_1 = 0,5 + \Phi(z_1),$$

звідки випливає, що потрібно знайти за додатком таке значення z_1 , для якого

$$\Phi(z_1) = q_1 - 0,5,$$

після чого визначається

$$x_1 = \mu + z_1 \sigma.$$

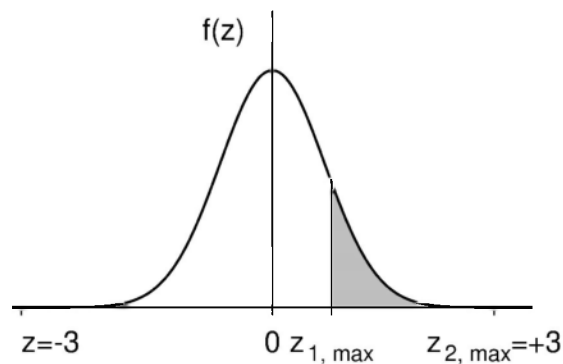


Рис. 6.17. Максимальний розв'язок оберненої задачі для $q_1 \geq 0,5$, $z_1 \geq 0$

Отримані максимальні розв'язки для початкової випадкової величини X та стандартизованої Z позначимо відповідно як $(x_{1 \max}, x_{2 \max})$ та $(z_{1 \max}, z_{2 \max})$.

Два побудованих розв'язки, мінімальний і максимальний, повністю визначають множину всіх можливих розв'язків у такому розумінні. Нехай із яких-небудь міркувань вибирається значення z_1 для шуканого розв'язку, при цьому має місце нерівність

$$z_{1 \min} \leq z_1 \leq z_{1 \max}.$$

Тоді z_2 можна визначити так, як і при побудові мінімального розв'язку:

$$z_{2 \min} \leq z_2 \leq z_{2 \max}.$$

Нехай певним чином вибирається значення z_2 для шуканого розв'язку, причому очевидно, що

$$z_{2 \min} \leq z_2 \leq z_{2 \max}.$$

Тоді z_1 можна визначити так, як і при побудові максимального розв'язку:

$$z_{1 \min} \leq z_1 \leq z_{1 \max}.$$

$$\text{Значення функції Лангаса} \quad \Phi(z) = \frac{1}{\sqrt{2\pi}} \int_0^z e^{-\frac{t^2}{2}} dt$$

Таблиця 6.4

z	0,00	0,01	0,02	0,03	0,04	0,05	0,06	0,07	0,08	0,09
0,0	0,00000	0,00399	0,00798	0,01197	0,01595	0,01994	0,02392	0,02790	0,03188	0,03586
0,1	0,03983	0,04380	0,04776	0,05172	0,05567	0,05962	0,06356	0,06749	0,07142	0,07535
0,2	0,07926	0,08317	0,08706	0,09095	0,09483	0,09871	0,10257	0,10642	0,11026	0,11409
0,3	0,11791	0,12172	0,12552	0,12930	0,13307	0,13683	0,14058	0,14431	0,14803	0,15173
0,4	0,15542	0,15910	0,16276	0,16640	0,17003	0,17364	0,17724	0,18082	0,18439	0,18793
0,5	0,19146	0,19497	0,19847	0,20194	0,20540	0,20884	0,21226	0,21566	0,21904	0,22240
0,6	0,22575	0,22907	0,23237	0,23565	0,23891	0,24215	0,24537	0,24857	0,25175	0,25490
0,7	0,25804	0,26115	0,26424	0,26730	0,27035	0,27337	0,27637	0,27935	0,28230	0,28524
0,8	0,28814	0,29103	0,29389	0,29673	0,29955	0,30234	0,30511	0,30785	0,31057	0,31327
0,9	0,31594	0,31859	0,32121	0,32381	0,32639	0,32894	0,33147	0,33398	0,33646	0,33891
1,0	0,34134	0,34375	0,34614	0,34850	0,35083	0,35314	0,35543	0,35769	0,35993	0,36214
1,1	0,36433	0,36650	0,36864	0,37076	0,37286	0,37493	0,37698	0,37900	0,38100	0,38298
1,2	0,38493	0,38686	0,38877	0,39065	0,39251	0,39435	0,39617	0,39796	0,39973	0,40147
1,3	0,40320	0,40490	0,40658	0,40824	0,40988	0,41149	0,41309	0,41466	0,41621	0,41774
1,4	0,41924	0,42073	0,42220	0,42364	0,42507	0,42647	0,42786	0,42922	0,43056	0,43189
1,5	0,43319	0,43448	0,43574	0,43699	0,43822	0,43943	0,44062	0,44179	0,44295	0,44408
1,6	0,44520	0,44630	0,44738	0,44845	0,44950	0,45053	0,45154	0,45254	0,45352	0,45449
1,7	0,45543	0,45637	0,45728	0,45818	0,45907	0,45994	0,46080	0,46164	0,46246	0,46327
1,8	0,46407	0,46485	0,46562	0,46638	0,46712	0,46784	0,46856	0,46926	0,46995	0,47062
1,9	0,47128	0,47193	0,47257	0,47320	0,47381	0,47441	0,47500	0,47558	0,47615	0,47670
2,0	0,47725	0,47778	0,47831	0,47882	0,47932	0,47982	0,48030	0,48077	0,48124	0,48169

Список рекомендованой литературы

- Алимов Ю.И. Является ли вероятность “нормальной” физической величиной? // Успехи физических наук. – 1992. – Т. 162, № 7. – С. 149-182; <http://www.ufn.ru>.
- Архипов С. Еще раз о законе Бенфорда // Техника – молодежи. – 1980. – № 8. – С. 61.
- Афифи А. Статистический анализ: Подход с использованием ЭВМ: Пер. с англ. / А.Афифи, С.Ейзен. – М.: Мир, 1982. – 488 с.
- Биометрия / Н.В. Глотов, Л.А. Животовский, Н.В. Хованов, Н.Н.Хромов-Борисов. – Л.: Изд-во ЛГУ, 1982. – 263 с.
- Борщ В.Л. Збірник задач з математичної статистики (для студентів біолого-медичного інституту) / В.Л. Борщ, Ю.П. Совіт. – Д.: ДДУ, 1997. – 80 с.
- Борщ В.Л. Математичні методи в біології в прикладах і задачах / В.Л. Борщ, Ю.П.Совіт. – Д.: ДДУ, 1998. – 80 с.
- Вентцель Е.С. Теория вероятностей / Е.С. Вентцель, В.И. Овчаров. – М.: Наука, 1970. – 348 с.
- Владимирский Б.М. Математические методы в биологии.– Ростов: Изд-во Рост. ун-та, 1983. – 304 с.
- Гильдерман Ю.И. Закон и случай. – Новосибирск: Наука, 1991. – 200 с.
- Гильдерман Ю.И. Лекции по высшей математике для биологов. – М.: Наука, 1974. – 410 с.
- Гласс Дж. Статистические методы в педагогике и психологии: Пер. с англ. / Дж. Гласс, Дж. Стэнли. – М.: Прогресс, 1976. – 495 с.
- Глотов Н.В. Сборник задач по биометрии / Н.В. Глотов, А.А. Филатов, Н.Н. Хромов – Борисов – Л.: Изд-во ЛГУ, 1985. – 81 с.
- Джонсон Н. Статистика и планирование эксперимента в технике и науке. Методы обработки данных: Пер. с англ. / Н. Джонсон, Ф. Лион – М.: Мир, 1980. – 462с.
- Дубровский В.Н. Как возникает распределение Пуассона // Квант. – 1988. №8. – С. 23 – 25.
- Евсеев Л. Закон Бенфорда // Техника – молодежи. – 1979. – № 10. – С. 53.
- Жалдак М.І. Теорія ймовірностей і математична статистика з елементами інформаційної технології / М.І. Жалдак, Н.М. Кузьміна, С.Ю. Берлінська. – К.: Вища шк., 1995. – 351 с.
- Каминский Л.С. Медицинская и демографическая статистика. – М.: Статистика, 1973. – 351 с.
- Лакин Г.Ф. Биометрия. – М.: Высш. шк., 1980. – 293 с.
- Математическая статистика /В.Н. Иванова, В.Н. Калинина, Л.А. Нешумова, И.С. Решетникова. – М.: Высш. шк., 1981. – 327 с.
- Мизес Р. Вероятность и статистика. – М.; Л.: Госиздат, 1930. – 250 с.
- Пригожин И. / И. Пригожин, И. Стенгерс. Порядок из хаоса. Новый диалог человека с природой. – М.: Прогресс, 1986. – 434 с.
- Плохинский Н.А. Биометрия. – М.: Изд-во МГУ, 1970. – 367 с.
- Рокицкий П.Ф. Биологическая статистика. – Мн.: Вышэйш. шк., 1973. – 320 с.

Секей А. Парадоксы статистики: Пер. с англ. – М.: Мир, 1974. – 297с.

Снедекор Дж.У. Статистические методы в применении к исследованиям в сельском хозяйстве и биологии. – М.: Изд-во с.-х. лит., 1961. – 503 с.

Суходольский Т.Р. Основы математической статистики для психологов. – Л.: Изд-во ЛГУ, 1972. – 429 с.

Терентьев П.В. Практикум по биометрии / П.В. Терентьев, Н.С. Ростова. – Л.: Изд-во ЛГУ, 1977. – 152 с.

Тутубалин В.Н. Вероятность, компьютеры и обработка результатов эксперимента // Успехи физ. наук. – 1993. – Т. 163, № 7. – С. 93 – 109; <http://www.ufn.ru>.

Урбах В.Ю. Статистический анализ в биологических и медицинских исследованиях. – М.: Медицина, 1975. – 294 с.

Фишер Р.А. Статистические методы для исследователей. – М.: Госстатиздат. – 1958. – 268 с.

Хинчин А.Я. Учение Мизеса о вероятностях и принципы физической статистики // Успехи физ. наук. – 1929. – Т. 9, вып. 2. – С. 141-166; <http://www.ufn.ru>.

Чукова Ю.П. Распределение Пуассона // Квант. – 1988. – № 8. – С. 15-18.

Martin P. Probability as a Physical Motive // Entropy. – 2007. V. 9, P. 42-57

Зміст

Передмова	3
1. Елементи теорії множин	6
1.1. Поняття множини	6
1.2. Відношення між множинами	7
1.3. Дії над множинами	8
1.4. Властивості дій над множинами	14
2. Елементи комбінаторики	15
2.1. Основні поняття комбінаторики	15
2.2. Принцип математичної індукції	15
2.3. Правило добутку для множин	15
2.4. Перестановки	19
2.5. Розміщення	20
2.6. Сполучення	21
3. Випадкові події	23
3.1. Причинність і випадковість у природознавстві	23
3.2. Поняття досліду та події	25
3.3. Ймовірність випадкової події	26
3.4. Статистична ймовірність	27
3.5. Класична ймовірність	31
3.6. Порівняння статистичної та класичної ймовірностей	33
3.7. Геометрична ймовірність	34
3.8. Простір елементарних подій	35
3.9. Дії над подіями	36
3.10. Правило додавання ймовірностей для двох несумісних подій	37
3.11. Правило додавання ймовірностей для декількох несумісних подій	39
3.12. Правило додавання ймовірностей для двох довільних подій	39
3.13. Умовна ймовірність	42
3.14. Правило множення ймовірностей для двох довільних подій	46
3.15. Правило множення ймовірностей для декількох довільних подій	46
3.16. Незалежні події	49
3.17. Правило множення ймовірностей для двох незалежних подій	51
3.18. Правило множення ймовірностей для декількох незалежних подій	52
3.19. Двохетапні дослідження	52
4. Випадкові величини	56
4.1. Поняття випадкової величини	56
4.2. Позначення, пов'язані з випадковими величинами	57
4.3. Класифікація випадкових величин	57
4.4. Опис випадкової величини	59

4.5. Поняття генеральної та вибіркової сукупностей	60
4.6. Класифікація вибірок (Методи витягнення вибірок)	61
4.7. Повний опис дискретної випадкової величини	63
4.8. Повний опис неперервної випадкової величини	68
4.9. Порівняння понять функції густини та фізичної “густини”	79
4.10. Функція розподілу випадкової величини	83
4.11. Вибіркові числові характеристики	87
4.12. Властивості вибірових середнього значення та дисперсії	92
4.13. Генеральні числові характеристики	93
4.14. Зв'язок вибірових і генеральних числових характеристик	94
4.15. Властивості генеральних середнього значення та дисперсії	94
4.16. Опис випадкової величини	95
4.17. Квантілі неперервної випадкової величини	100
4.18. Побудова моделі випадкової величини за її описом	104
5. Стандартні розподіли дискретних випадкових величин	106
5.1. Біномний розподіл	106
5.2. Розподіл Пуассона	111
5.3. Розподіл Бенфорда	114
6. Стандартні розподіли неперервних випадкових величин	116
6.1. Рівномірний розподіл	116
6.2. Показниковий розподіл	118
6.3. Нормальний розподіл	119
6.4. Функція Лапласа	125
Список рекомендованої літератури	136

Темплан 2007, поз. 21

Навчальне видання

Володимир Леонідович Борщ
Юрій Петрович Совіт

Математичні методи в біології.
Основи теорії ймовірностей
Конспект лекцій

Редактор О.В.Бец
Техредактор Л.П.Замятіна
Коректор Т.А.Андреєва

Підписано до друку 05.11.07. Формат 60×84/16. Папір друкарський.
Друк плоский. Ум. друк. арк. 8,13 Ум. фарбовідб. 8,13 Обл.-вид. арк. 7,57
Тираж 200 пр. Зам. № .

РВВ ДНУ, вул. Наукова, 13, м. Дніпропетровськ, 49050.
Друкарня ДНУ, вул. Наукова, 5, м. Дніпропетровськ, 49050