

МИНИСТЕРСТВО ОБРАЗОВАНИЯ И НАУКИ УКРАИНЫ
НАРОДНАЯ УКРАИНСКАЯ АКАДЕМИЯ

И. С. Нечитайло, М. В. Бирюкова

МАТЕМАТИЧЕСКИЕ МЕТОДЫ В СОЦИОЛОГИИ

Учебник
для студентов высших учебных заведений

Харьков
Издательство НУА
2013

УДК 316:[51:371] (075.8)
ББК 60.5в631я73-1
Н59

*Рекомендовано
Министерством образования и науки Украины
в качестве учебника для студентов высших учебных заведений
Письмо № 1/11-323 от 16.01.2013*

Рецензенты:

д-р физ.-мат. наук *А. А. Янцевич* (Харьковский национальный университет им. В. Н. Каразина);
д-р социол. наук *Н. Н. Саппа* (Учебно-научный институт права и массовых коммуникаций Харьковского национального университета внутренних дел);
д-р социол. наук *В. Л. Погребная* (Национальный университет «Юридическая академия Украины имени Ярослава Мудрого»).

Підручник складається з п'яти розділів, в яких висвітлено особливості застосування математичних методів в соціології. Матеріали підручника, дозволять студентам закріпити відповідні теоретичні знання, удосконалити навички застосування математики в практичній діяльності соціолога, розвинути компетенції бакалавра соціології.

Підручник призначений для студентів вищих навчальних закладів, які навчаються за спеціальністю «Соціологія» та цікавляться питаннями математичної обробки соціологічної інформації.

Нечитайло, Ирина Сергеевна.

Н59

Математические методы в социологии : учебник для студентов высш. учеб. заведений / И. С. Нечитайло, М. В. Бирюкова ; Нар. укр. акад., [каф. социологии]. – Харьков : Изд-во НУА, 2013. – 320 с.

Учебник состоит из пяти разделов, в которых освещены особенности применения математических методов в социологии. Материалы учебника позволят студентам закрепить соответствующие теоретические знания, усовершенствовать навыки применения математики в практической деятельности социолога, развить компетенции бакалавра социологии.

Учебник предназначен для студентов высших учебных заведений, которые обучаются по специальности «Социология» и интересуются вопросами математической обработки социологической информации.

УДК 316:[51:371] (075.8)
ББК 60.5в631я73-1

ВВЕДЕНИЕ

Одна из важнейших задач, стоящих перед преподавателем, который читает курс, связанный с изучением математики для студентов-гуманитариев, обществоведов (психологов, социологов, юристов и др.), – это задача убедить аудиторию студентов в том, что математика им, действительно, необходима. Причем решать эту задачу необходимо на первом же занятии.

Нам в руки не попадалось такого учебника «по математике» для студентов нематематических специальностей, который не начинался бы со вступительной речи-обращения автора к студентам, с теми или иными доводами в пользу важности для них математических знаний. Проанализировав всю совокупность этих доводов, приведем те из них, которые, на наш взгляд, представляются наиболее убедительными.

Во-первых, математика способствует развитию абстрактного мышления как ни одна другая наука. Математика изучает модели реальных процессов, которые описываются на математическом языке. Человек, владеющий этим языком, способен глубже проникнуть в суть реальных процессов, совершенно не связанных с математикой, правильно ориентироваться в окружающей действительности [1, с. 3]. Человек, формулирующий математическое утверждение, проводящий математическое доказательство, оперирует истинно научным знанием, а также не обыденной, а предметной речью, строящейся по определенным законам: краткость, четкость, лаконичность, минимизация и т. д. Математическое образование следует рассматривать как важнейшую составляющую фундаментальной подготовки бакалавра. Обусловлено это тем, что математика является не только мощным средством решения прикладных задач и универсальным языком науки, но также и элементом общей культуры.

Современная математика становится, так сказать, «междисциплинарным инструментарием», который выполняет две основ-

ные функции: 1) обучающую (умение правильно задавать цель тому или иному процессу, определять условия и ограничения в достижении этой цели), 2) аналитическую, то есть «проигрывания» на искусственных моделях возможных в реальности ситуаций, видения траекторий развертывания этих ситуаций и путей разрешения проблем, которые могут возникнуть.

Следовательно, посредством изучения математики происходит развитие абстрактного мышления, способности к абстрагированию и умения работать с абстрактными, воображаемыми, неосознаваемыми объектами, без которых не обходится ни одна наука. Кроме того, в процессе усвоения основных правил осуществления математических процедур формируется ассоциативное и алгоритмическое мышление, развивается его сила, гибкость, конструктивность. Эти качества мышления сами по себе могут быть и не связаны с каким-либо конкретным тематическим содержанием и с математикой вообще. Однако понимание математики вносит в их формирование важную и специфическую, неотъемлемую и необходимую составляющую, которая не может быть внесена ни одной другой дисциплиной.

И еще одна важная особенность, присущая именно математике: она воспитывает такой склад ума, который требует критической проверки и логического обоснования тех или иных положений и точек зрения. Элемент сомнения – это здоровое рациональное зерно, присущее процессу математического мышления – нигде и никогда не помешает любому профессионалу [1, с. 5].

Во-вторых, математика способствует развитию ключевых компетенций. Не раскрывая подробно смысл этого понятия, скажем, что в общем ключевые компетенции опираются на универсальные знания, умения, на обобщенный опыт творческой деятельности и т. п. [15, с. 34–42]. Приведем перечень некото-

¹ **Андрей Викторович Хуторской** (р. 6 октября 1959 г.) – доктор педагогических наук, член-корреспондент Российской академии образования, академик Международной педагогической академии, директор Института образования человека, директор Центра дистанционного образования «Эйдос».

рых их них, которые выделяются А. В. Хуторским¹: 1) ценностно-смысловая; 2) общекультурная; 3) учебно-познавательная; 4) информационная; 5) коммуникативная; 6) социально-трудовая; 7) личностная (самосовершенствование) [27].

Очевидно, что ключевые компетенции являются важным аспектом междисциплинарного подхода в образовании. Перво-степенной задачей для преподавателя, читающего математический курс на факультете, где обучаются «не-математики», является выработка эффективных методов, способов и средств формирования ключевых компетенций посредством изучения математики именно у них – у студентов-гуманитариев [31, с. 5].

Что же касается студентов, специализирующихся именно в области социологии, то, пожалуй, согласимся с высказыванием одного из выдающихся российских ученых Г. Г. Татаровой² о том, что у них наблюдается «синдром аллергии на формулы, графики и т. д.». Вряд ли найдется преподаватель, читающий математический курс на социологическом факультете, который захочет оспаривать это высказывание. Вместе с тем очевиден тот факт, что именно математическая грамотность является неотъемлемой составной частью профессиональной подготовки каждого специалиста, имеющего дело с анализом массовых явлений, будь то социальные, экономические и др. Социологическое изучение любого сложного объекта, имеющего социальную природу, не может ограничиваться чистым теоретизированием. В противном случае социология как отдельная наука об обществе рискует полностью слиться с социальной философией. Проникновение математических методов в разнообразные сферы человеческой деятельности предоставляет социологии весьма эффективные средства для исследования сложных социальных объектов. Для того чтобы результаты подобных исследований могли считаться достоверными, а их выводы обоснованными, обязательно следует прибегать к определенному математическому формализму. Без такого формализма современному социологу просто не обойтись,

² **Татарова Галина Галеевна** – доктор социологических наук, профессор, почетный доктор Института социологии РАН.

так как насущная необходимость получения адекватной, однозначной и максимально точной информации социологического характера систематически возрастает. Все чаще приходится сталкиваться с сомнениями относительно результатов тех или иных социологических исследований. И, к сожалению, следует признать, что такие сомнения отнюдь не всегда безосновательны. Проблема в том, что многие социологические исследования носят далеко не научный характер и проводятся на низком профессиональном уровне, что наиболее часто связано с неумением владеть математическими методами работы с эмпирическими данными и нежеланием социолога обращаться к глубокому математическому анализу этих данных (работа – не из легких). Однако для обеспечения высокого качества результатов исследовательской деятельности социолога математические знания и умения представляются крайне необходимыми. Такая необходимость диктуется как динамикой социальных трансформаций во всем мире, так и нынешним этапом развития социальных процессов в нашей стране. Стабильность и устойчивость общества зависит от того, насколько предсказуемы социальные процессы и явления. Современные задачи социального управления, прогнозирования и планирования невозможно решать без глубоких знаний в области обработки и анализа социологической информации с использованием устойчивых математических методов.

Именно поэтому в настоящее время от работника, занятого в любой области, связанной с изучением массовых социальных явлений, в том числе и от социолога, требуется высокий уровень математической грамотности. В связи с этим для студентов, обучающихся социологическим специальностям, большое значение имеет ознакомление с общими категориями, принципами и методологией математического анализа, усвоение ряда математических методов, представляющих важность в обеспечении высокого качества работы социолога. Для этого в обязательную образовательную программу вузов вводятся специальные дисциплины, позволяющие получить соответствующие знания, умения и навыки, развить важные компетенции. Одной из таких

дисциплин является дисциплина «Математические методы в социологии».

Одноименный учебник содержит полный объем информации, необходимой для всестороннего и углубленного анализа эмпирических данных, полученных в результате социологического исследования, с применением математических методов. В своей структуре он имеет пять разделов. *Раздел I* посвящен детальному рассмотрению понятия измерения в социологии. *Раздел II* акцентирует внимание на методах агрегирования данных и числовых характеристиках исследуемой совокупности, таких как среднее арифметическое, мода, медиана, дисперсия, среднее квадратическое отклонение, коэффициент вариации и др. Кроме того, во втором разделе рассмотрены основные правила табличного и графического представления социологических данных. *Раздел III* ставит во главу угла вопросы, связанные с формулировкой, а также подтверждением/опровержением статистических гипотез и характеристикой ошибок, которые могут возникнуть при анализе социологических данных. В *Разделе IV* особое внимание уделяется вопросам, связанным с определением связи между признаками, особенностями вычисления коэффициентов, которые могут свидетельствовать о наличии/отсутствии этой связи, ее силе (тесноте) и направлении. И, наконец, *Раздел V* посвящен рассмотрению тех математических методов, к которым социолог может прибегнуть при анализе данных исследований, проведенных в малых социальных группах (на примере социометрического опроса).

Следовательно, учебник «Математические методы в социологии» содержит тот оптимум информации, который необходим студенту, обучающемуся по социологической специальности для приобретения знаний относительно теоретических законов распределения генеральной и выборочной совокупности, навыков вычисления количественных характеристик социальных объектов, интерпретации полученных данных.

Надеемся, что структура и содержание материала, изложенного в учебнике, будет способствовать формированию целостного видения логики исследовательского процесса, связанного

с анализом социологических данных и использованием возможностей, предоставляемых математическим аппаратом для обработки и анализа этих данных. Цели и задачи изучения особенностей применения математических методов в социологии направлены на формирование системы компетенций, необходимых для полноценной социально-профессиональной деятельности выпускника высшего учебного заведения, его самореализации. В частности, речь идет о таких компетенциях, как:

- социально-личностные (организация личной деятельности, которая отображается умением учитывать изменения, происходящие в обществе, адаптироваться к ним и др.);
- общенаучные (умение использовать в социальных и профессиональных практиках базовые знания в области современных информационных технологий; навыки использования программных средств и др.);
- общепрофессиональные (обеспечение высокого качества осуществляемой профессиональной деятельности социолога, что отображается умением качественно осуществлять экспертизу данных, полученных в ходе социологического исследования);
- специализированно-профессиональные (обеспечение необходимого уровня индивидуальной подготовки в области профессиональных знаний, использования собственных знаний и наблюдений, а также сложных процедур социологического анализа).

Раздел I. ИЗМЕРЕНИЕ КАК ПРЕДПОСЫЛКА ПРИМЕНЕНИЯ МАТЕМАТИЧЕСКИХ МЕТОДОВ В СОЦИОЛОГИИ

1.1. Измерения в социологии: теоретико-методологический ракурс

Что понимается под измерением в социологии? Как известно, измерение и подсчет являются основой математики как фундаментальной науки, результаты изысканий которой находят широкое применение практически во всех областях научного знания. Математика активно используется в целях достижения истины представителями и естественных, и гуманитарных наук. Если использование математики в естественных науках является традиционным и считается само собой разумеющимся, то целесообразность ее связи с гуманитарными науками не является такой очевидной. Тем не менее, все гуманитарные и социальные науки стремятся обладать аппаратом, позволяющим осуществить эмпирическую проверку выдвигаемых гипотез с выходом на точные, конкретные, обоснованные выводы. Не является исключением из ряда этих наук и социология.

В одном из своих учебников академик РАН Г. В. Осипов, рассуждая о приоритетных направлениях применения математики в социологии, выделяет следующие: построение и расчет выборки, анализ данных, моделирование, измерение. Однако в итоге рассуждений делается заключительный вывод о том, что наиболее важной среди проблем применения математики в социологии является социологическое измерение [2, с. 22–23]. Соглашаясь с этим выводом, считаем необходимым более подробно остановиться на рассмотрении проблемы измерения в социологии.

Обозначенная в заголовке тема акцентирует внимание на двух основных моментах: во-первых, на уточнении самого понятия

«измерение», а во-вторых, на выяснении того, «что?» и «как?» измеряется в социологии.

Итак, что же такое измерение? Однозначного ответа на данный вопрос нет и, в общем-то, никогда не существовало. А в условиях междисциплинарности, мультипарадигмальности, поликонтекстуальности и т. п. современного научного знания перспективы поиска такой однозначности в отношении трактовки данного термина (и не только его) представляются не только весьма отдаленными, но и не совсем целесообразными [31]. Тем не менее, для формирования единого, целостного видения роли математических методов в социологии (что является залогом их дальнейшего успешного применения на практике) необходимо все-таки попытаться сформулировать то определение измерения, которое наиболее гармонично вписывается в систему социологического знания и соответствует, так сказать, «социологическому мировоззрению».

Анализ всех существующих определений привел ученых, занимающихся данной проблемой, к пониманию того, что традиционно противопоставляются два основных подхода к пониманию измерения [2, с. 34–39]. Первый из них может быть условно назван *дескриптивным*, поскольку задачей измерения здесь является описание, определение существующей величины. Представители этого подхода, традиции которого берут начало в античной науке и наиболее активно развиваются намного позже *Б. Расселом*³, настаивают на том, что измерить можно лишь те объекты, которые априори обладают конкретной величиной. Например, рост человека может быть большим, маленьким, средним и т. п. Его величина – совершенно объективная мера, которая может быть представлена в виде числа и не зависит ни от желания того, кто измеряет, ни от желания того, кого измеряют. А измеряемый объект заведомо обладает определенными числовыми свойствами.

Представители второго подхода, конструктивного, настаи-

³ **Бертран Артур Уильям Рассел** (1872–1970) – английский математик, философ и общественный деятель.

вают на том, что ни один объект не обладает априорной величиной, а его числовые свойства создаются, конструируются в процессе измерения. Основные положения данного подхода наиболее ярко представлены в трудах известного английского специалиста в области проблем измерения *Н. Кемпбэлла*⁴. Исходя из этих положений, такой объект измерения, как человеческий рост, не содержит в себе никаких величин и чисел. Он может быть наделен совершенно разными числами, а сам числовой показатель роста зависит от того, кто измеряет, кого измеряют, с какой целью и в каких условиях. К примеру, рост 156 см – это низкий рост для взрослого человека, но высокий рост для восьмилетнего ребенка. Если обладателю такого роста этот показатель кажется слишком маленьким, то он будет склонен считать и говорить, что его рост около 160 см. В таком случае прибавленные 4 см не играют роли, однако в случае приема на военную службу, к примеру, скрупулезно подсчитывается каждый сантиметр. Если рост человека измеряется в Украине, то ему приписывается число 156 (см), а если этот же человек поедет, например, в США, то это число изменится коренным образом и будет представлено такими числами: 5' 1", что означает 5 футов и 1 дюйм.

Если вдуматься, то спор между представителями двух этих подходов сродни классическому противостоянию «объективистов» и «субъективистов» в социологии. Однако вспомним тот факт, что в современной социологии особой популярностью пользуются теории, объединяющие объективное и субъективное по принципу дополнительности, а не тотального отвержения, в духе диалектического синтеза. Примером могут стать такие из них, как, например, теория социального пространства и теория габитуса П. Бурдьё, теория структуризации Э. Гидденса, теория коммуникативного действия Ю. Хабермаса и социетальных систем Н. Лумана и др.

Если разобраться, то и в отношении определения понятия

⁴ **Кемпбэлл Норман Роберт** (1880–1949) – английский физик, философ, работавший в области философии науки.

«измерение» вполне возможно избежать каких-либо противопоставлений. В самых общих чертах **измерение** можно представить как процесс, совокупность операций, направленных на определение числовых характеристик объекта, которые являются информативно важными в отношении описания данного объекта и объяснении процессов и явлений, связанных с этим объектом. Для того чтобы сгладить противостояние обозначенных выше подходов, можно некоторые измерения обозначить как прямые, (например, измерение длины проградуированной линейкой) и косвенные (основанные на известной зависимости между искомой величиной и непосредственно измеряемыми величинами) [12]. Измерение выполняется при помощи специальных измерительных средств, с целью нахождения числового значения измеряемой величины в принятых единицах измерения.

Весь комплекс проблем, которые возникают в связи с измерением, мы назовем проблемой измерения. Проблема измерения, как и всякая научная проблема, включает в себя, прежде всего, три аспекта – философско-гносеологический, теоретико-методологический и практический. На эту тему можно философствовать очень долго. Однако наш интерес должен быть сосредоточен на вопросах, которые связаны с проблемами измерения в социальных науках в целом, а в частности, в социологии.

Если разобраться, то социологии как отдельной науки вообще могло бы не существовать, если бы великие праотцы социологической мысли не увидели определенные закономерности в функционировании и развитии общества, не осознали необходимость и возможность подтверждения существования этих закономерностей с помощью эмпирических исследований. В основе любого эмпирического исследования, как правило, лежит измерение исследуемого объекта. Начнем с поиска ответа на вопросы: «что?» и «зачем?» измеряется в социологии. Для того чтобы в этом разобраться, следует вспомнить, что вообще изучают и исследуют социологи. Как известно, предметом социологии, является изучение общества или его отдельных фрагментов (составляющих единиц, элементов, процессов, явлений) сквозь призму их социальной организации, социальных

связей и отношений, одним словом – социального взаимодействия.

Согласно одному из социологических словарей, внимание социологов должно быть сфокусировано на изучении социального взаимодействия его различных форм и проявлений на всех его уровнях (макро-, мезо-, микро)⁵, с целью анализа структуры социальных отношений, складывающейся в ходе этого взаимодействия [9, с. 390].

Однако чисто теоретических рассуждений для достижения этой цели, как мы уже говорили, недостаточно. В рамках социологического знания развивается перспективное направление – прикладная социология, каждый элемент которой конкретизирует общее представление о предмете науки, методах исследования, способах его организации и оформления полученных результатов. Основная функция прикладной социологии видится в возможности научного обеспечения решения задач социологии на практике, разрешении реально существующих проблем, связанных с социальным взаимодействием в различных сферах общества. Прикладная социология направлена на установление и констатацию научных фактов, создающих базис социологического знания. Научный факт – это зафиксированное посредством научного метода событие, явление, которое используется для подтверждения выводов. На основании научных фактов определяются свойства и закономерности, выводятся теории и законы. Переход от фактов к теории предполагает несколько промежуточных ступеней. Первой такой ступенью является описание факта. Дело в том, что научный факт не определяется ни личными знаниями исследователя, ни его органами чувств, ни какими-либо другими особенностями ученого как отдельной личности-субъекта. Объективность факта связана со спецификой изменения предмета познания в эксперименте. Измерительные процедуры являются важнейшим, а чаще всего, необходимым

⁵ **Макроуровень** – взаимодействие на уровне социальных систем, подсистем и институтов; **мезоуровень** – взаимодействие на уровне социальных групп и организаций; **микроуровень** – взаимодействие отдельных индивидов.

и обязательным дополнением процедуры эксперимента. Именно на основе установленных количественных характеристик в науке часто удается сделать выводы относительно строения объекта познания и законов его функционирования.

Огюст Конт, создавая методологическую базу социологии, в числе основных ее методов определил наблюдение социальных фактов. Следуя таким представлениям, можно сказать, что любое социологическое исследование, носящее прикладной характер, является своеобразным «экспериментом», в ходе которого, устанавливаются, обосновываются, подтверждаются или опровергаются те или иные социальные факты. На этой основе прикладная социология разрешает задачи преобразования, прогнозирования и планирования социальных процессов. Эти задачи неразрывно связаны с математической обработкой социологической информации, что и создает «научно обоснованный» базис для разработки планов и выработки решений, построения прогнозов и моделей развития того или иного социального объекта [5, с. 9].

Данные социологических исследований сами по себе не позволяют сделать обобщенные выводы, выявить тенденции, проверить гипотезы, так как основные понятия, используемые при конкретно-социологических исследованиях, это, как правило, абстрактные понятия. Измерить их в буквальном смысле этого слова с помощью традиционных измерительных средств, приборов, установок и т. п. (например, линейки, мерной ложки или амперметра) – нереально. Величины, характеризующие социальные явления, в ряде случаев не подчиняются аксиомам арифметики. Мы с точностью можем сказать, на сколько числовых единиц, во сколько раз «10» больше «5». Однако мы не можем с арифметической точностью сказать и выразить ответ в конкретных числовых единицах, определяя, насколько статус одного человека или группы людей выше/ниже статуса остальных. Тем не менее, мы можем иметь целый ряд оснований, чтобы сделать такой вывод.

По словам одного из создателей современной теории измерений К. Кумбса: «...Измерение в физических науках обычно

означает приписывание чисел наблюдениям (этот процесс называется отображением), и анализ данных состоит в операции с этими числами. Социальный ученый, беря физику за образец, часто пытается делать то же самое. Существует мнение, что социальный ученый, который следует этой процедуре, иногда разрушает свои данные» [2, с. 35–39].

В социологии измерение – процедура, при помощи которой объекты исследования, рассматриваемые как носители определенных социальных отношений и составляющие эмпирическую систему, отображаются в некую метрическую систему с соответствующими отношениями между ее элементами. В качестве объектов социологического измерения могут выступать индивиды, группы индивидов (трудовой, студенческий коллектив), их социальные качества, условия, в которых они существуют и взаимодействуют, и многое другое. При измерении в социологии, как правило, объектам приписываются определенные элементы, используемые в качестве эталона математических систем (элементами здесь являются действительные числа, отношения между которыми в процессе измерения могут быть различными). Однако возможно использование и нечисловых систем: частично упорядоченных множеств, графов, матриц и т. д. Объекты измерения вступают в отношения как носители определенных свойств. Поэтому вместо словосочетания «измерение объектов» часто используется – «измерение свойств объектов» [15].

Тот подход к пониманию измерения, который находит наиболее широкое практическое применение в социологии в настоящее время, начал формироваться на рубеже XIX–XX веков. Его возникновение было обусловлено острой потребностью общественных наук, которые именно к этому времени достигли того уровня, когда дальнейшее интенсивное их развитие без использования формальных моделей изучаемых процессов или явлений стало немыслимым. Непригодность классического подхода к измерению в общественных науках обусловила расширение этого понятия, вследствие чего под измерением стал пониматься способ приписывания чисел объектам независимо от того,

использовалась ли при этом единица измерения. В основе такого подхода лежит предположение о существовании изоморфизма или гомоморфизма (одинаковой устроенности, похожести между эмпирическими и числовыми системами). Очевидно, классическое понимание измерения не противоречит такому подходу и может рассматриваться как частный случай последнего [15].

Понятие измерения в социологии представляется настолько широким, что распространяется и на случаи использования нечисловых математических систем. Хотя в научной литературе можно встретить работы, в которых оспаривается право называть измерением те процедуры, в которых задействованы такие системы (шкалы), где не содержится никаких цифр (номинальные и некоторые порядковые), а в результате получаются очень интересные и, главное, – достоверные результаты, что подтверждено многолетней практикой.

Использование чисел в жизни и деятельности людей отнюдь не всегда предполагает, что эти числа можно складывать и умножать, производить иные арифметические действия. Ну, например, телефонные номера состоят из одних только цифр – отличная почва для всевозможных математических процедур! Однако большинство из этих процедур не имеют никакого смысла. Что можно сказать о человеке, который занимается умножением телефонных номеров? [9, с. 5]. Ведь, если речь идет об исследовании живых существ, а не тестировании машин, результаты могут быть совершенно непредсказуемыми. Скажем, если поместить в клетку двух животных, то через сутки в этой же клетке можно найти только одно животное (случай с кошкой и мышкой) либо больше двух (случай с овцой, которая окотилась в суточный период эксперимента). В любом случае, ни для кого не секрет, что область применения чисел намного шире, чем арифметика.

Методологические основы измерений в социологии. Любое событие или объект, окружающие человека, можно представить в виде системы компонент или совокупности составляющих его частей: химический состав человеческого тела, «рецепт» пирога, факторы, обуславливающие заболевание, социальные условия, вызывающие войну, или характеристики социального

института. Цель естественных и общественных наук заключается в расчленении явления с тем, чтобы сделать возможным его понимание. Однако недостаточно просто правильно назвать факторы, определяющие явление, нужно измерить их «удельный вес», интенсивность – количественную меру. Химическая формула, анализ крови, доза лекарства, интенсивность социальной установки, размер населения – все эти данные свидетельствуют о том, что измерение играет первостепенную роль в познании.

Сам факт измерения предполагает, что имеется единица измерения. Как в естественных, так и в общественных науках одна из основных проблем заключается в создании эффективных единиц измерения, которые позволяли бы проводить исследование в непрерывно изменяющемся мире. Естественные науки создали с этой целью лабораторный метод, который позволяет «остановить» явление, повторить его, изучить предмет под микроскопом или нагреть его до точки кипения, не говоря уже о более сложных способах наблюдения и измерения. А если говорить о человеке, обществе и таком, довольно абстрактном феномене, как социальное взаимодействие, то в этом случае данные собираются и квантифицируются с намного более значительными трудностями. Измерение социальных явлений встречает много серьезных препятствий. Ведь общество неоднородно и изменчиво, многие его характеристики не укладываются в стандартизованные понятия или категории. Социальную дистанцию, например, труднее измерить, чем географическое расстояние, социальные силы кажутся иллюзорными, по сравнению с физическими. Тем не менее основная гипотеза, которая служит поводом для использования математических методов в социологии, сводится к тому, что все окружающее человека существует в больших или меньших количествах и, как правило, повторяется с определенной периодичностью. Естественно, что в связи с этим у социологов возникает большое число резонных вопросов, на которые им самим же и придется искать ответы.

Ну, например, разве можно говорить об абсолютной точности числа представителей такой социальной группы, как молодежь, если с каждым днем численность данной группы изменяется.

Ситуация усложняется еще и тем, что само определение молодежи имеет множество вариаций. Возрастные границы данной группы, как правило, варьируются в пределах от 14–16 до 28–35 лет. А если вспомнить крылатую фразу: «Человеку столько лет, на сколько он себя чувствует», то определение каких-либо границ и вовсе представляется невозможным.

Как можно подтвердить измерением простое утверждение, что граждане Украины сейчас вступают в брак в более раннем возрасте, чем несколько десятков лет назад, что безработица также уменьшилась, что преступность возросла в послевоенные годы, а мужчины являются лучшими водителями, чем женщины, или что интерес к современной музыке возрастает? Как можно подсчитать, измерить факторы, влияющие на преступность подростков или на частоту разводов?

Все эти вопросы можно объединить двумя требованиями:

- 1) необходимостью объективного количественного ответа;
- 2) наличием соответствующих единиц измерения, которые далеко не очевидны и почти всегда находятся с большим трудом. В общем случае способ квантификации определяется характером рассматриваемых переменных.

Измерение в социологии – это способ изучения социальных явлений с помощью количественных оценок. В современном понятии измерение – это процедура, с помощью которой объекты социологического исследования отображаются на определенную числовую или графическую систему. Основными компонентами измерения при этом являются:

- 1) объекты измерения;
- 2) свойства объектов (т. к. измеряются не объекты, а их свойства);
- 3) шкалы, на которых отображаются измерения.

Методологическими основами измерений в социологии являются определение и обоснование объекта измерения, свойств, шкал, выборки, методов и инструментов сбора информации.

1.2. Совокупности, признаки, переменные: к вопросу о том: «что измеряется в социологии?»

Совокупности, признаки, переменные: соотношение понятий. В социологии одним из центральных и наиболее часто используемых понятий является понятие социальной группы. С социологической точки зрения социальная группа – совокупность людей, объединяемых на основании общих социально значимых признаков, то есть таких, которые связаны с их вовлеченностью в социальное взаимодействие. Скажем, цвет глаз может рассматриваться как объединяющий признак, но он не является социально значимым, так как не обуславливает специфику социального взаимодействия всех голубоглазых, кареглазых и др. А вот материальное положение является признаком социально значимым, так как людей, объединяемых на основании данного признака, отличает не только толщина кошелька, но и общность тех или иных форм социального взаимодействия (трудовая и политическая активность, например).

В прикладной социологии социальная группа также рассматривается как совокупность, но сам термин «совокупность», как правило, употребляется в сочетании с такими словами, как «генеральная» либо «выборочная», что подчеркивает его эмпирическую специфику.

По словарному определению, *генеральная совокупность* – это множество элементов, обладающих каким-то одним или многими признаками.

Выборочная совокупность – это конкретное множество объектов (субъектов), отобранных специальным образом для обследования (опроса). Исследование, проведенное на основе выборочной совокупности, называется выборочным. Выборочный метод позволяет экономить временные и материальные затраты на него. К примеру, вместо двух миллионов, представляющих генеральную совокупность, мы можем опросить около тысячи человек и распространить полученные выводы на всю совокупность. Такое распространение будет вполне объективным, однако лишь при том условии, что выборочная совокуп-

ность репрезентативна, то есть основные ее характеристики совпадают с характеристиками генеральной совокупности.

Совокупность изучается социологом на основании присущих всем ее объектам общих признаков. *Признак* можно определить как исследуемую особенность генеральной совокупности. Важным является то, что каждый из них может подвергаться измерению.

Для того чтобы осуществить полное, всестороннее исследование того или иного социального процесса, явления и др., безусловно, чистого теоретизирования недостаточно. Необходимо каким-то образом осуществить переход от теоретических понятий исследования к эмпирически интерпретируемым понятиям, а затем – к эмпирическим индикаторам. Другими словами, обязательно необходим поиск «чего-то» такого, за чем можно понаблюдать, а результаты наблюдения выразить конкретным значением. Речь идет об *эмпирической операционализации* основных понятий как о процессе перевода основных теоретических положений в форму, удобную для эмпирической проверки, то есть форму эмпирических индикаторов.

«*Эмпирический индикатор*» – понятие, употребляемое в прикладной социологии не менее часто, чем понятие генеральной или выборочной совокупности. По сути, это доступная наблюдению и измерению характеристика (признак) изучаемого или управляемого социального объекта, то есть доступный наблюдению признак.

Признак можно представить как величину, которая изменяется от объекта к объекту. Следовательно, признак – *переменная* величина. Например, все люди обладают таким признаком, как материальное положение. Однако у каждого человека «свое» материальное положение: кто-то имеет много денег и может ни в чем себе не отказывать, а у кого-то денег не хватает даже на приобретение вещей первой необходимости. В целом, вариаций материального положения, как и вариаций преобладающего большинства других признаков, с помощью которых можно охарактеризовать совокупность, состоящую из людей, может быть довольно много. Исследователю необходимо постараться

охватить вниманием как можно больше всех возможных вариаций.

Признаки, как переменные, могут быть классифицированы. Классической является дифференциация переменных на:

- 1) качественные/количественные;
- 2) дискретные/непрерывные.

Качественные и количественные переменные. Одни признаки отражают качественные характеристики элемента исследуемой совокупности (или всего объекта исследования), а другие – количественные. Следовательно, каждый признак может изменяться, варьироваться либо качественно, либо количественно.

Качественные (фиктивные) переменные – это переменные, которые не могут иметь количественных числовых значений. Значение количественной переменной выражается числом, значение качественной – текстовым описанием, рисунком или каким-либо другим поясняющим его смысл способом.

Качественные переменные могут изменяться от наблюдения к наблюдению, но не численно, а скорее по характерным отличиям. Эти характерные отличия обычно называются «атрибутами». Кстати, атрибуты могут быть расположены в определенном порядке, в зависимости от интенсивности присутствия определенного качества у объекта исследования, на основе чего становится возможным упорядочение объектов, их сравнение в каком-то отношении друг другом. Например, такой, на первый взгляд, количественный признак, как возраст, может быть представлен в виде качественной переменной и варьироваться следующим образом: молодой, зрелый, пожилой, старый и т. п., а о таком признаке, как вес нам может говорить такое перечисление качеств объекта, как: худой, среднего телосложения, полный, толстый и т. п. Каждый, кто понимает смысл этих обозначений, будет понимать их количественно. И наоборот, такие очевидно качественные признаки, как религиозные взгляды, профессиональная принадлежность и др., можно подвергнуть ранжированию. Например, религиозные взгляды могут рассматриваться как более или менее ортодоксальные,

а профессии могут быть иерархизированы по их социальному престижу.

Некоторые авторы употребляют термин «вариация» только по отношению к тем переменным, которые могут изменяться количественно. А качественные данные иногда определяются как «неизменные» или «не выразимые в числах», чем подчеркивается только то, что они не могут быть выражены в числах, а лишь перечисляются как характерные черты. Количество и качество, согласно таким определениям, были бы полярными категориями. Тем не менее в последнее время все больше появляется высказываний в пользу возможности применения термина «измерение» как к качественным, так и к количественным данным.

К *количественным переменным* относятся те переменные, которые выражаются в численных величинах, то есть существуют в больших или меньших количествах. Примерами количественных переменных могут служить: возраст, рост, доход, размер населения, размер семьи, срок заключения, данные о рождаемости и т. д. Люди характеризуются большим или меньшим возрастом, города – различной численностью населения, осужденные – различными сроками заключения. Эти переменные могут быть измерены, и каждое число, получающееся в результате измерения, называется значением переменной. Совокупность таких значений может быть упорядочена от наименьшего значения к наибольшему. Подобная возможность жесткой упорядоченности может служить удобным признаком, отличающим количественные переменные от качественных.

В свою очередь, количественные переменные могут подразделяться на дискретные и непрерывные. Рассмотрим критерии такой дифференциации более подробно.

Дискретные и непрерывные переменные. *Дискретной* (прерывистой, состоящей из отдельных частей) считается та переменная, которая может принимать значения только из некоторого списка определенных чисел. Дискретные переменные представлены перечнем вариаций, выраженных числами «неделимыми», с точки зрения логики и здравого смысла. Примерами дискретной переменной являются число детей в семье (не может

быть 2,5 ребенка); число совершаемых преступлений (невозможно совершить 3,2 преступления) и т. п. Значения дискретной переменной, опять-таки, исходя из здравого смысла, невозможно выстроить в непрерывный ряд, в непрерывную линию. Если измерять совокупность по признаку количества детей в семье, то расстояние между одним ребенком, двумя, тремя и более нелогично заполнять значениями 1,1; 1,2; 1,3 ...2,8; 2,9 и т. д., так как в природе таких показателей не существует.

В отличие от дискретных, *непрерывные* переменные могут делиться на меньшие части, теоретически – до бесконечности. Возраст, расстояние, вес, могут принимать любое из бесчисленного множества значений. Если упорядочить все это множество по возрастанию или убыванию, можно легко получить неделимую, непрерывную линию. Если, к примеру, говорить о таком признаке-переменной, как возраст, то все возможные его вариации, значения можно дробить бесконечно. Фактически, можно определить возраст каждого человека с точностью до долей секунд, обозначив его конкретной цифрой. Чтобы прожить один год своей жизни, человек должен пройти через каждый миг этого временного континуума.

Измерение непрерывной никогда не может быть абсолютно точным; всегда остается возможность еще более точного измерения. За делением в одну тысячную миллиметра на микрометре, одну долю секунды на хронометре существуют еще меньшие деления. Это означает, что непрерывные данные всегда являются округленными.

Несмотря на то что обычно непрерывные переменные рассматриваются как округленные, а дискретные как измеренные точно, необходимо более осторожно использовать термин «точное измерение».

На практике как дискретные, так и непрерывные переменные определяются только приближенно. Дискретный счет точен только тогда, когда сама шкала счета достаточно мала и когда объект изменяется достаточно медленно. Так, например, отец может объявить, что имеет восемь детей – вполне точная мера в течение определенного времени. Однако в большинстве

социологических исследований кажущийся точный счет – рождения, разводы, размер населения из-за изменчивости самого этого социального явления и технической невозможности получить исчерпывающие данные по существу является приближенным.

Часто в качестве примера дискретной переменной приводится денежная система, так как копейку, например, нельзя разделить на меньшие доли. С точки зрения здравого смысла, это не подлежит сомнению. Однако на практике любая непрерывная переменная не может дробиться неограниченно. Не следует смешивать теоретическую возможность с практически доступной. Эти рассуждения могут показаться несколько педантичными, но необходимо отметить, что различие между непрерывными и дискретными переменными играет первостепенную роль в различных статистических моделях.

Если дана генеральная совокупность лиц, которые изучаются, например, по своему доходу, то в этом случае признак «доход» может быть выражен количественной непрерывной величиной, принимающей в определенных пределах любые значения. Если же эти N лиц изучаются на предмет их семейного положения по признаку размера семьи, в этом случае варианта является величиной дискретной (прерывной), поскольку она может принимать только лишь значения целых чисел: 1, 2, 3... и т. д.

1.3. Шкалы, типы шкал, возможные операции со шкалами: к вопросу о том: «как происходит измерение в социологии?»

Для того, чтобы осуществить любое, даже самое элементарное измерение с большой вероятностью точности, необходимо владеть специальным измерительным прибором или средством. Как известно, разнообразие таких приборов и средств настолько велико, что их перечисление может занять не одну страницу текста. Общим же для всех них является обязательное наличие шкалы.

Самым широко используемым инструментарием социологического измерения является анкета (опросный лист, бланк

интервью и т. п.). Как известно, ни один из анкетных вопросов не включается в нее «просто так». Каждый из них тщательно продумывается, проходит через апробацию. Вопрос анкеты нацелен на то, чтобы, получив ответы на него, исследователь смог сделать конкретные выводы в отношении интересующего признака исследуемой совокупности. И если для заурядного респондента вопрос анкеты – это только вопрос и не более того, то для исследователя-социолога – это, прежде всего, шкала, с помощью которой измеряется тот или иной признак исследуемой совокупности как варьирующаяся переменная.

Следовательно, во главе угла обозначенной темы поставим поиск ответов на вопросы: «Что такое шкала?», «Какие бывают шкалы в социологии?» и «Как работают эти шкалы?», то есть какие измерения они позволяют осуществить и какую ценную информацию социолог может получить в результате. Если ответы на эти вопросы будут найдены, можем считать, что основная цель данного параграфа выполнена.

Шкалы, с помощью которых социолог измеряет признаки объектов исследуемой совокупности, применяемые социологом, такое же необходимое и незаменимое средство, как шкала линейки закройщика для портного или шкала тонометра для врача.

В общем понимании, *шкала* представляет собою некую знаковую, систему, для которой задано гомоморфное отображение. Говоря простым языком, эта система представляет собой совокупность отметок (точек, штрихов и др. условных обозначений), расположенных в определенной последовательности. Рядом с некоторыми из этих отметок проставлены числа или другие символы, соответствующие ряду последовательных значений измеряемой величины [12].

Различные типы измерительных шкал широко используются в физике, экономике, психологии, социологии и др. Особый случай измерения требует подбора особого измерительного средства, прибора, аппарата и т. п. Одним из основоположников такого подхода к пониманию измерения стал американский психолог С. С. Стивенс, автор общеизвестной классификации шкал по уровню измерения. Впоследствии его идеи были развиты

П. Суппесом, Дж. Зинесом, Р. Д. Льюсом и др., что привело к активному использованию в эмпирической социологии шкал как эталонов и средств измерения [21].

Вспомним, что для социолога измерение – это процедура, с помощью которой объекты социологического исследования отображаются в определенной числовой и/или графической системе. Такими системами являются шкалы. Шкала – совокупность всех возможных значений признака (которые могут изменяться от объекта к объекту). Представление социологической информации в виде шкал открывает возможности осуществления с признаком, имеющим социальную природу, ряда математических операций, что в итоге позволяет исследователю сделать конкретные выводы в отношении объекта изучения.

В этой связи возникает вполне резонный вопрос, связанный выбором соответствующих исследовательским целям тех или иных широко известных соотношений между числами, чтобы анализ шкальных значений позволил нам сделать содержательные выводы. Для ответа на этот вопрос необходимо в первую очередь четко представить себе характер числовых систем, использующихся в процессе измерения в социологии [17, с. 138].

Измерить – значит сравнить с эталоном. Применение числовых методов измерения и повлекло появление шкал. Именно шкала служит эталоном измерения. Следуя наиболее широко используемому определению, **шкала** – это измерительная часть инструмента, оценивающая эмпирические индикаторы, в том числе с позиции их расположения в определенной последовательности. Процесс построения шкалы, проведенный по определенным правилам, называют шкалированием.

При проведении измерений в социологии необходимо осуществить выбор шкалы. Шкалы дифференцируются по типам. Выбор типа шкалы зависит от того, какой в виде переменной представляется признак. Например, качественные переменные могут быть измерены с помощью классификационной шкалы (номинальной), когда в результате измерения определяется принадлежность объекта к определенному классу или с помощью порядковой шкалы, когда в результате измерения также опреде-

ляется не только принадлежность объекта к некоторому классу, но и осуществляется определенное упорядочение объектов, их сравнение друг с другом в каком-то отношении. Например, диагноз заболевания – это измерение в классификационной шкале, а определение степени тяжести заболевания – измерение в порядковой шкале

Но, как это обычно бывает в науке, не существует одной единственной типологии шкал – их, по меньшей мере, несколько. Первыми разработчиками типологии шкал для измерения социальных признаков считаются ученые, работающие в области теории измерения Дж. Зинес и П. Суппес. Несколько упрощенные, рассчитанные на социологов, интерпретации этой типологии можно найти в многочисленных работах отечественных и зарубежных ученых, посвященных методологии и методам социологического исследования, а также возможностям применения математических процедур при анализе данных социологических исследований (В. И. Паниотто, А. Б. Телейко, Ю. Н. Толстовой, Г. Г. Татаровой, В. А. Ядова и др.).

Конструирование того или иного типа шкалы зависит от того, какой именно признак и с какой целью мы хотим измерить, в виде какой переменной хотим его представить. При построении шкалы необходимо учитывать, что:

- шкала должна отражать именно те свойства исследуемой совокупности, которые измеряются, и учитывать все их значения – валидность;
- шкала должна быть достаточно чувствительной (речь идет о достаточном количестве градаций);
- желательно, чтобы все позиции оценок на шкале располагались симметрично (значение с «+» симметрично значению с «-»);
- необходимы достаточная точность и надежность шкалы (устойчивость к изменению объекта), чтобы объективно отразить картину измерений, учтя все возможные варианты значений.

Шкалы принято классифицировать по типам измеряемых данных, которые определяют допустимые для данной шкалы математические преобразования, а также типы отношений,

отображаемых соответствующей шкалой. Современная классификация шкал была предложена в 1946 году американским психологом Стенли Смитом Стивенсом⁶.

Рассмотрим для начала дифференциацию шкал по принципу *низкого/высокого* типа. Например, Ю. Н. Толстова, описывая различия между этими шкалами пишет, что *шкалы низкого типа* – это шкалы, позволяющие получать «числа», очень не похожие на те действительные числа, с которыми мы имеем дело в математике. Эта «непохожесть» означает невозможность работать с этими числами по обычным правилам арифметики. Давая определение *шкалам высокого типа*, ученый подчеркивает, что к ним причисляют те, с помощью которых получаются числа, в достаточной мере похожие на действительные числа, то есть такие, с которыми мы можем сделать почти все, что мы привыкли делать с числами [26].

Шкалами низкого типа обычно считают шкалы, называемые в литературе номинальными и порядковыми (состоящими из слов), а к шкалам высокого типа относят шкалы метрические (состоящие из цифр).

Рассмотрим более подробно эти шкалы, их характеристики и математические процедуры, которые возможно осуществить с данными, полученными в результате измерения с помощью этих шкал.

Шкалы низкого типа. Итак, *номинальная шкала* (классификационная, шкала наименований) – шкала низкого типа, которая используется для измерения значений качественных признаков и *допускает отношение равенства-неравенства*. Например, числами, которые, как правило, приписываются игрокам командного вида спорта, мы не можем оперировать, как с числами в арифметике, а именно два игрока первого номера нам не дадут игрока второго номера и т. п. [2, с. 33]. Значением

⁶ **Стивенс (Stevens) Стэнли Смит** (р. 4.11.1906, Огден, штат Юта), американский психолог. Профессор Гарвардского университета (с 1946), руководитель лабораторией психофизики (организованной им же в 1944).

такого признака является наименование класса эквивалентности, к которому принадлежит рассматриваемый объект. Следовательно, номинальная шкала устанавливает отношения равенства между объектами, входящими в одну категорию. Каждой категории приписывается число. Примером шкалы наименований могут служить: различия по полу, национальности, ценностным ориентациям и т. д. Измерение с помощью номинальной шкалы – самый простой вид измерения.

Признаки, измеряемые с помощью номинальной шкалы, удовлетворяют аксиомам тождества [13]:

- либо $A = B$, либо $A \neq B$;
- если $A = B$, то $B = A$;
- если $A = B$ и $B = C$, то $A = C$.

Для данной шкалы вычисляют следующие числовые характеристики: частоты распределения; моду; находят критерий согласия Хи-квадрат, критерий Крамера (для проверки гипотезы о связи качественных признаков) и др.

Как говорилось выше, помимо номинальных шкал к шкалам низкого типа относят и **порядковые** (ранговые) **шкалы**. Они строятся на отношении тождества (как и номинальные), а также порядка, то есть *допускают отношение равенства-неравенства и больше-меньше*. Возможные значения признака, заданные в градациях порядковой шкалы, ранжированы. Такая шкала груба, потому что не учитывает разности между градациями. Пример такой шкалы – оценки успеваемости, выраженные словами: «неудовлетворительно», «удовлетворительно», «хорошо», «отлично». С помощью порядковых шкал часто измеряются установки респондентов в отношении кого-либо или чего-либо. Например, если измеряется признак отношения к реформам в образовании, то градации соответствующей порядковой шкалы могут определяться следующими словами с присвоением каждому из них следующих рангов: «положительно» (+2), «скорее положительно» (+1), «нейтрально» (0), «скорее отрицательно» (–1), «отрицательно» (–2).

Для данной шкалы вычисляют следующие числовые характеристики: все те же, что и для номинальной шкалы, а также

медиану; квартили; коэффициент Спирмена; коэффициент Кенделла; коэффициент проверки социологических гипотез.

Для количественных признаков применяются **метрические шкалы**, которые считаются **шкалами высокого типа** и подразделяются на интервальные шкалы и шкалы отношений.

Интервальная шкала используется только для числовых значений признака. Она устанавливает не только порядок между объектами и их свойствами, но и разность или расстояние между объектами, то есть *допускает отношение равенства-неравенства, больше-меньше и равенства-неравенства интервалов*. По отношению к этой шкале возможно осуществление основных арифметических операций. Для нее определяются следующие числовые характеристики: все те же, что для номинальной и порядковой шкал, а также среднее арифметическое (простое и взвешенное); дисперсия, среднее квадратическое отклонение; коэффициенты корреляции (частной и множественной); коэффициент регрессии.

Интервальная шкала довольно редко применяется в социологических исследованиях, так как если речь идет об измерении социальных явлений, процессов, феноменов, то установить равенство между объектами или указать точное расстояние между ними крайне тяжело.

Шкала отношений (абсолютная шкала) – это метрическая интервальная шкала, в которой присутствует дополнительное свойство: естественное и однозначное присутствие нулевой точки. Потому ее и называют абсолютной, то есть она является единственной из четырех шкал, имеющей абсолютный ноль. Нулевая точка характеризует полное отсутствие измеряемого качества. Определение нулевой точки – сложная задача для социологических исследований, накладывающая ограничение на использование данной шкалы [19]. Классическими примерами шкалы отношений являются шкалы пространственных мер (длина, высота, ширина), массы, денежных единиц и др. Шкалой отношений данная шкала называется потому, что она показывает не только конкретные числовые величины, но и отношение между этими величинами, то есть отношение «во столько-то раз

больше». Примером может служить измерение признака численности аудитории. Шкала отношений при этом будет включать следующие (или аналогичные) градации: «до 5 человек», «5–10 человек», «10–15 человек», «15 и более человек» и т. д. Значения шкалы отношений можно умножать на константу. Она *допускает операции равенства-неравенства, больше-меньше, равенство интервалов и равенство отношений* и тем самым реализует все арифметические операции, то есть к перечисленным для всех предыдущих шкал добавляются вычисления средней гармонической и средней геометрической [2, с. 30].

В социологической литературе можно встретить случаи, когда шкалы низкого типа (и получаемые с их помощью данные) определяют как *качественные*, а шкалы высокого типа (и соответствующие данные) – как *количественные*, или *числовые*. Однако на сегодняшний день среди социологов-практиков, профессионально занимающихся проведением эмпирических исследований, больше противников, чем сторонников такого определения. Основная проблема измерения в социологии связана с тем, что интересующие социолога данные, как правило, бывают получены именно по шкалам низких типов, что существенно усложняет применение в социологии традиционных математических методов. Учитывая эти ограничения, может показаться, что статистический метод неэффективен при анализе качественных переменных. Тем не менее в настоящее время со стороны ученых ведется упорная работа по постоянному совершенствованию методов исследования и обработки качественных переменных. Эти методы имеют очень большое значение для социологов, поскольку многие важные социологические данные могут быть исключительно качественными по своей форме.

Существуют сложные варианты измерения при помощи других шкал. Одной из них является шкала Богардуса. Основное предназначение этой шкалы – измерение национальных и расовых установок. Ее особенность заключается в том, что каждая оценка (мнение, позиция) автоматически включает в себя все последующие и исключает все предыдущие. Вопрос, который используется при измерении подобной шкалой, имеет следую-

щую формулировку: «Какие взаимоотношения с представителем «такой-то» национальности для вас приемлемы?»: брачные отношения; личная дружба; соседство; быть коллегами по работе; быть жителями одного города, поселка, села и т. д. Опыт свидетельствует о том, что подобные шкалы могут быть построены и успешно использоваться для измерения установок относительно явлений в различных сферах социального взаимодействия.

Шкалы в социологии могут быть дифференцированы и типологизированы по уровню измерения, в зависимости от следующих характеристик: качество, порядок, расстояние и наличие начальной точки (см. *Табл. 1.1*).

Таблица 1.1

**Типы шкал по уровню измерения
и их основные характеристики**

Уровень измерения	Характеристики шкал			
	качество	порядок	расстояние	наличие начальной точки
Шкала наименований	*			
Шкала порядка	*	*		
Интервальная шкала	*	*	*	
Шкала отношений	*	*	*	*

Некоторые методы измерения (методы построения шкал). Числа, полученные в результате применения «традиционных» для социологического исследования шкал, в одних случаях могут служить непосредственной оценкой измеряемого качества, а в других – основой для дальнейшей математической обработки и построения производной шкалы.

Построение шкал методом экспертных оценок. Измерение в этом случае разбивается на два этапа: построение шкалы, то есть построение шкальных весов признаков, и оценивание респондентов по этим шкалам.

Метод равных интервалов. При большом числе признаков метод парных сравнений оказывается громоздким, поскольку эксперты должны рассмотреть каждую возможную пару

признаков, а число таких пар быстро растет с ростом числа признаков ($k(k-1)/2$).

Суть метода равных интервалов состоит в следующем: список суждений оценивается экспертами, которые располагают признаки в фиксированное число категорий (обычно равное 7, 9 или 11), ранжированных по степени предпочтения. Интервалы между категориями должны быть субъективно равными, то есть судящему должно казаться, что различия по степени предпочтения между суждениями любых двух смежных категорий равны.

Из списка отбираются 100–200 суждений, которые по возможности равномерно покрывают весь континуум установки: от мнений, выражающих крайнее положительное отношение к объекту установки, до мнений, выражающих крайнее отрицательное отношение к этому объекту. Особое внимание уделяется формулировке нейтральных мнений. При отборе суждений необходимо руководствоваться следующими правилами:

- суждения формулируются в утвердительной форме и должны выражать сиюминутную психологическую установку испытуемого, не смешивая ее с отношением к тому же объекту в прошлом;
- каждое суждение должно быть достаточно кратким, чтобы не утомлять опрашиваемых;
- суждения должны быть сформулированы таким образом, чтобы их можно было принять или отвергнуть, и выразить они должны одну, а не несколько идей; необходимо избегать утверждений, по отношению к которым большинство респондентов не имеют определенного мнения;
- суждения должны иметь такую форму, чтобы согласие или несогласие с ними указало на их отношение к объекту установки;
- характер суждений не должен быть фактическим;
- необходимо исключить утверждения, которые могут быть приняты как теми, кто имеет отрицательное отношение к объекту установки, так и теми, кто имеет положительную установку.

Следующая процедура состоит в отборе судей (от 50 до 300 человек). Экспертов просят независимо друг от друга разложить карточки на 7, 9 или 11 групп так, чтобы в последней

группе были карточки с утверждениями, соответствующими максимально положительной установке индивида, высказавшего эти утверждения; в первую группу помещались карточки с утверждениями, соответствующими максимально негативной установке к объекту исследования; в среднюю группу помещались карточки с нейтральными утверждениями.

После сортировки необходимо оценить каждое суждение с точки зрения его соответствия шкале и установить вес суждения на шкале. На этом этапе построения шкалы суждению, помещенному данным экспертом в некоторую категорию, приписывается число (оценка), совпадающая с номером этой категории. Затем вычисляется медиана распределения оценок данных всей группы экспертов для каждого суждения.

Присуждение баллов респонденту. Когда с помощью одного из рассмотренных методов шкала построена, дальнейшая процедура заключается в том, чтобы присвоить каждому респонденту в массовом опросе балл. Для этого отобранные суждения в беспорядке тасуются и заносятся в опросный лист. Веса суждений при этом не указываются.

Обычно задание респонденту состоит в том, что его просят отметить те суждения, с которыми он согласен. При этом естественно ожидать, что он метит ограниченное число суждений. За количественную оценку респондента берутся медианы шкальных значений отмеченных им суждений.

Метод суммарных оценок. Предложен в 1932 г. Р. Лайкертом. Идея довольно проста: группе лиц даются вопросы, которые должны оцениваться по пятибалльной системе в отношении согласия с этими вопросами.

Баллы одного лица относительно всех вопросов суммируются. Полученная сумма – балл этого лица. Затем лица ранжируются по баллам.

Для построения шкалы отбирается большое число утверждений, относящихся к исследуемой установке.

Шкалограммный анализ. Шкалограммный анализ Гутмана ведет к построению шкал порядкового уровня измерения. Эта

техника связана с построением одномерных шкал, то есть шкал, не затрагивающих вопросов или не включающих факторов, посторонних по отношению к измеряемой характеристике.

Основная идея метода состоит в том, что шкала должна состоять из иерархизированной системы вопросов, то есть такой, в которой согласие с вышестоящим по иерархии суждением должно вести к согласию с нижестоящими суждениями.

Процедура построение шкалы предполагает ряд этапов:

1. Отбирается серия суждений относительно измеряемого свойства.

2. Эти суждения раздаются группе респондентов (около 100 человек), в которую входят представители обследуемой категории населения. Респонденты отвечают на каждый вопрос либо «да», либо «нет».

3. Отбрасываются те суждения, которые набрали более 80% благожелательных и отрицательных ответов. Число оставшихся признаков должно быть не менее десяти.

4. Следующий шаг состоит в ранжировании оставшихся вопросов и респондентов по числу набранных баллов от высшего к низшему. В идеальном случае должна получиться картина (шкалограмма).

Затем вычисляется коэффициент воспроизводимости:

$$\text{Воспроизводимость} = 1 - \frac{\text{число ошибок}}{\text{число ответов}}$$

Если воспроизводимость не менее 0,9, – это означает, что данный набор суждений образует одномерную шкалу.

Кроме всего прочего, необходимо учитывать следующие критерии:

- каждая категория (суждение) должна обладать минимальной ошибкой;
- ошибки должны иметь случайный характер.

Если какая-то одна частная ошибка встречается значительно чаще, чем другая, – это значит, что признак не принадлежит шкальному типу. Суждение, которое не удовлетворяет этим требованиям, отбрасывается.

Семантический дифференциал. Метод семантического дифференциала (СД) разработан Чарльзом Осгудом⁷ для измерения смысла понятий и слов и, прежде всего, для дифференциации эмоциональной стороны значения данного понятия, для изучения эмоциональных компонентов социальных установок.

Методы измерения. Поскольку наблюдаемые переменные не являются однородными по своей природе, постольку очевидно, что ни один из методов измерения в отдельности не будет полностью эффективен при обработке всех видов переменных. Можно назвать в порядке объективности и надежности четыре метода измерения:

- 1) прямое перечисление;
- 2) использование абстрактных стандартных единиц измерения;
- 3) моделирование наблюдаемого явления;
- 4) построение упорядочивающих шкал.

Прямое перечисление является простым подсчетом признаков без использования новой терминологии. Признак, сформулированный в определенных категориях, сам является единицей измерения. Таким образом, число разводов, преступлений людей или число правильных и неправильных ответов на вопросы – все это прямое перечисление определенных единиц. Каждый из перечисляемых признаков считается «равным самому себе», возможные вариации и различия данного признака в перечисляемых объектах игнорируются. Это простейший тип квантификации, доступный любому мальчишке, пересчитывающему свои марки.

Как указывалось выше, прямому перечислению иногда не хватает необходимой точности, так как оно пренебрегает различиями и вариациями признаков. Бананы, которые, вообще говоря, различаются по размерам, могут продаваться поштучно (прямое перечисление) или на вес, килограммами. Переводя бананы в килограммы, используется на несколько более высоком уровне абстракции *стандартная единица измерения*, которая является

⁷ **Чарльз Осгуд** (англ. *Charles Osgood*), 20 ноября 1916 – 15 сентября 1991 – американский психолог.

эталоном сравнения для исследуемых объектов. Килограмм, сантиметр, год – все это условные нормы измерения, понятные каждому человеку. Они заменили более примитивные эталоны сравнения, такие как «шаг», «локоть» и «лошадиная сила». Некоторые из этих эталонов еще живут в фольклоре, но все они уже перестали удовлетворять стандартам точности, принятым в современной промышленности и торговле.

Каждый из этих двух типов измерения дает объективные данные, которые легко могут быть повторены любым другим наблюдателем.

Третий метод измерения использует то, что называется эквивалентом повеления, позволяющим проводить косвенные наблюдения, менее надежные, чем в двух вышеописанных методах. В общественных науках многие материалы исследования представляют различные типы установок и ценностей, имеющих исключительно важное значение для понимания человеческого поведения. Однако эти установки недоступны непосредственному наблюдению, их нельзя подсчитать или измерить, приложив линейку. Общественное мнение, социальные установки, социальный престиж, семейное счастье, расовая вражда, групповая мораль – каждое из этих социальных явлений представляет интерес для социолога, однако их субъективный, психологический характер не позволяет применять вышеописанные непосредственные методы измерения. Поэтому вынуждены обращаться к явлениям, которые могут быть непосредственно изучены и измерены как операциональный эквивалент субъективных переменных. Такая ускользающая величина, как семейное счастье, может быть измерена с помощью таких эквивалентов, как продолжительность брака, совместная деятельность супругов, число конфликтов и т. д. Изменения в музыкальном вкусе можно проследить по содержанию репертуара концертных исполнителей или по объему продажи различных аудиокассет, компакт-дисков.

Исследователи в области общественных наук проявили много изобретательности, конструируя эквиваленты для измерения непосредственно не воспринимаемых факторов социальной жизни.

Однако каждое косвенное измерение сопряжено со множеством проблем, от которых совершенно свободно непосредственное измерение. Прежде всего, очевидно, что «эквивалент» не эквивалентен явлению. Даже комбинация нескольких эквивалентов не могла бы идеально копировать объект, который она по предложению представляет.

Наконец, измерение может заключаться в *ранжировании* ряда объектов согласно выбранному критерию. Такие цепи *упорядочивания* могут быть сконструированы на основе уже описанных ранее объективных мер или могут быть чисто субъективными с более или менее явным обоснованием критерия предпочтения. Например, города могут быть упорядочены относительно объективно в соответствии с количеством проживающего в них населения, супружеские пары могут быть ранжированы по «счастью» после анализа соответствующих тестов. В данных примерах в процессе ранжирования приходится пренебрегать точностью, важно лишь установление порядка следования объектов. Примером субъективного ранжирования может служить определение учеником личного предпочтения по отношению к своим одноклассникам или ранжирования кинофильмов, производимого экспертами, которые, без сомнения, будут руководствоваться своими личными вкусами, более или менее осознанными. Точно так же могут быть ранжированы в порядке престижа профессии, соседи, поэты и композиторы. Вообще говоря, любой ряд признаков может быть ранжирован в соответствии с каким-либо внешним критерием, а любой ряд значений может быть упорядочен в соответствии с их численной величиной.

Очевидно, что технические неудобства ранжирования заключаются в присущей ему малой точности. Вследствие того, что шкала упорядочивания строится на основе чьих-то суждений или мнений: 1) интервалы между последовательными порядками или рангами на шкале не обязательно должны быть равными; 2) несколько ранжирований одних и тех же объектов не обязательно должны быть идентичными. Однако в тех случаях, когда возможны только субъективные нормы оценки или когда

точность не играет существенной роли, упорядочивание по рангам оказывается полезным.

Порядковые ранги, которые по существу являются просто индикаторами относительного положения на шкале, тем не менее часто могут служить численными весовыми коэффициентами, хотя и искусственными. Например, степени интенсивности установки часто взвешиваются весами от одного до пяти, чтобы соответствующие данные несли больше информации для последующего анализа. Такие условные аппроксимации весьма полезны и часто используются в социологических исследованиях.

Каждый из четырех вышеописанных методов измерения обладает своими индивидуальными характеристиками; каждому свойственны свои правила и принципы интерпретации. Более того, данное измерение может оказаться сочетанием двух или более методов.

Например, прямое перечисление преступлений или разводов, требуемое практикой судебных учреждений, может одновременно быть мерой социальной дезорганизации или указывать на непрочность семейных связей. Все эти методы, однако, объединяются одним общим и важным признаком – тем, что они представляют собой попытку выразить измерение в терминах и символах, которые делают возможным развитие и коммуникацию социального знания, закладывают основы техники измерения, без которой не существовало бы статистики.

С примерами вопросов социологической анкеты, сформулированных в различных шкалах измерений, можно ознакомиться в «Приложении 1», которое находится на последних страницах данного учебника.

1.4. Надежность измерения

Помимо описанных сложностей с пониманием сути понятия «измерение» в социологии, существует и целый ряд проблем, связанных с практическим обеспечением того, чтобы осуществленная процедура измерения соответствовала этой сути, то есть, действительно, считалась настоящим измерением. Такое соответствие достигается за счет высокой надежности измерения.

В процессе измерения участвуют три составляющие: объект измерения, измеряющие средства, с помощью которых производится отображение свойств объекта на числовую систему, и субъект, производящий измерение. Предпосылки надежного измерения кроются в каждой отдельной составляющей.

Следует отметить, что наиболее детально методы и техника контроля данных на надежность изложены в работах Г. И. Саганенко⁸ и В. И. Паниотто⁹ [18, с. 74–75; 23; 24]. Последний применяет аналитический подход к предмету, выделяя множество разновидностей надежности и технических приемов оценки ее уровня, тогда как Г. И. Саганенко акцентирует внимание на наиболее существенных, неперенных требованиях и сравнительно простых способах контроля надежности.

Не останавливаясь здесь на дискуссии терминологического характера, заметим, что в строгом смысле словосочетание «надежность измерения» правомерно относить именно к инструменту, с помощью которого производится измерение, но не к самим данным, подлежащим измерению. В отношении данных, как и заключительных выводов исследования, правильное говорить, что они достоверны (или относительно достоверны) в том числе и потому, что фиксированы надежным инструментом.

Надежность измерения следует понимать как «воспроизводимость» результатов измерения, повторяемого при идентичных условиях.

Каждое измерение подвержено ошибкам. Все ошибки измерения можно разделить на три класса: *промахи* (грубые ошибки), *систематические* и *несистематические* (случайные) ошибки.

Промахи. В процессе измерения иногда возникают грубые ошибки, причиной которых могут быть неправильные записи

⁸ Саганенко Галина Иосифовна – ведущий научный сотрудник филиала Института социологии РАН в Санкт-Петербурге.

⁹ Паниотто Владимир Ильич (р. 22 января 1947 г. в Киеве) – украинский социолог, доктор философских наук, генеральный директор Киевского международного института социологии, президент компании InMind, профессор кафедры социологии Университета «Киево-Могилянская Академия».

исходных данных, плохие расчеты, неквалифицированное использование измерительных средств и т. п. Это проявляется в том, что в рядах измерений попадают данные, резко отличающиеся от совокупности всех остальных значений. Чтобы выяснить, нужно ли эти значения признать грубыми ошибками, устанавливают критическую границу так, чтобы вероятность превышения ее крайними значениями была достаточно малой и соответствовала некоторому уровню значимости.

В социологическом исследовании «промах» в измерении признака может быть связан с личностными характеристиками исследователя, особенностями характера, уровнем профессионализма, мировоззренческой позицией, в конце концов, то есть иметь субъективный характер, а также с резким нарушением условий измерения при отдельных наблюдениях, то есть иметь объективные причины: резкое изменение социальных, экономических, политических и других условий (как следствие революции, экономического кризиса, экологической катастрофы и т. п.).

Систематические ошибки. Систематические ошибки – это ошибки, которые возникают из-за путаницы переменных в реальном мире или из-за особенностей самого инструмента. Они появляются каждый раз, когда применяется данный инструмент, и постоянно сопутствуют объектам и исследованиям, в которых используется одно и то же измерение. Постоянные ошибки делают наши результаты невалидными в том смысле, что различия (или сходства), которые, как представляется, выявляют наши измерения, не есть точные отражения различий, которые мы, по нашему мнению, измеряем.

Систематической ошибкой называют составляющую погрешности измерения, остающуюся постоянной и закономерно изменяющуюся при повторных измерениях одной и той же величины. Систематические ошибки наблюдаются, если, например, шкала линейки нанесена неточно (неравномерно); капилляр термометра в разных участках имеет разное сечение. Как видно из примеров, систематическая ошибка вызывается определенными причинами, величина ее остается постоянной (смещение нуля шкалы прибора, «неравноплечность» весов), либо

изменяется по определенному (иногда довольно сложному) закону (неравномерность шкалы, неравномерность сечения капилляра термометра и т. д.).

Случайные ошибки измерения. Случайной называют составляющую погрешности измерений, изменяющуюся случайным образом при повторных измерениях одной и той же величины. Случайные ошибки проявляются по-разному и обусловлены преходящими характеристиками объектов, ситуационными различиями в применении инструмента, ошибками в проведении измерения и обработке данных и другими факторами. Они делают наши измерения невалидными почти так же, как и систематические ошибки. Кроме того, случайные ошибки делают наши измерения *ненадежными* в том смысле, что проявление случайных ошибок не дает возможности постоянно получать одни и те же результаты при использовании одного и того же измерения.

При проведении с одинаковой тщательностью и в одинаковых условиях повторных измерений одной и той же постоянной неизменяющейся величины мы получаем результаты измерений – некоторые из них отличаются друг от друга, а некоторые совпадают. Такие расхождения в результатах измерений говорят о наличии в них случайных составляющих погрешности. Случайная ошибка возникает при одновременном воздействии многих источников, каждый из которых сам по себе оказывает незаметное влияние на результат измерения, но суммарное воздействие всех источников может оказаться достаточно сильным. Случайная ошибка может принимать различные по абсолютной величине значения, предсказать которые для данного акта измерения невозможно. Эта ошибка в равной степени может быть как положительной, так и отрицательной. Случайные ошибки всегда присутствуют в эксперименте. При отсутствии систематических ошибок они служат причиной разброса повторных измерений относительно истинного значения. Допустим, что при помощи секундомера измеряют период колебаний маятника, причем измерение многократно повторяют. Погрешности пуска и остановки секундомера, ошибка в величине отсчета, небольшая

неравномерность движения маятника – все это вызывает разброс результатов повторных измерений и поэтому может быть отнесено к категории случайных ошибок. Если других ошибок нет, то одни результаты окажутся несколько завышенными, а другие несколько заниженными. Но если, помимо этого, часы еще и отстают, то все результаты будут занижены. Это уже систематическая ошибка.

Проверка измерения на надежность предполагает поиск систематических и несистематических ошибок измерения, их соотнесение. Возможны различные типологии приемов оценки надежности первичной информации, например, с точки зрения внешнего или внутреннего контроля данных, получаемых определенным способом. Мы предлагаем пользоваться обобщающим понятием надежности инструмента измерения (и соответственно надежности данных, фиксируемых этим инструментом), имея в виду три составляющие: (1) обоснованность, (2) устойчивость и (3) правильность измерения. Естественно, что и методы контроля на надежность нужно рассматривать в этих трех аспектах.

Понятие *правильности* связано с возможностью учета в результате измерения различного рода систематических ошибок. При изучении правильности устанавливается общая приемлемость данного способа измерения. При этом систематические ошибки измерения могут проявляться в следующем.

- *Отсутствие разброса ответов по значениям шкалы.* Попадание ответов в один пункт свидетельствует о полной непригодности измерительного инструмента – шкалы. Такая ситуация может возникнуть или из-за «нормативного» давления в сторону общепринятого мнения, или из-за того, что градации (значения) шкалы не имеют отношения к определению данного свойства у рассматриваемых объектов [17, с. 143–145]. Например, если все опрашиваемые респонденты согласны с утверждением (хорошо, когда работа или задание требуют универсальных знаний), нет ни одного ответа (не согласен), остается только зафиксировать этот факт, однако подобная шкала не поможет дифференцировать изучаемую совокупность по отношению респондентов к работе.

- *Использование части шкалы.* Довольно часто обнаруживается, что практически работает лишь какая-то часть шкалы, какой-то один из ее полюсов с прилегающей более или менее обширной зоной. Так, если респондентам для оценки предлагается шкала, имеющая положительный и отрицательный полюса, в частности от + 3 до – 3, то при оценивании какой-то заведомо положительной ситуации респонденты не используют отрицательные оценки, а дифференцируют свое мнение лишь с помощью положительных. Для того чтобы вычислить значение относительной ошибки измерения, исследователь должен знать определенно, какой же метрикой пользуется респондент - всеми семью градациями шкалы или только четырьмя положительными. Так, ошибка измерения в один балл мало о чем говорит, если известно, какова действительная вариация мнений.

Например, девятнадцати респондентам было предложено выразить отношение к трем понятиям по семи шкалам по каждому. Шкалы имели по 21 градации с крайними полюсами +10 и –10 и средней точкой 0. В целом получено 399 (19*3*7) оценок со следующим распределением (см. *Табл. 1.2*):

Таблица 1.2

Распределение оценок респондентов

Балл (a_i)	10	9	8	7	6	5	4	3	2	1	0	-1	-2	-3	-4	-5	-6	-7	...	-10
Частота (n_i)	145	33	30	37	25	24	25	10	12	8	39	3			3			5		

Поскольку значения $a_i < 0$ использовались всего лишь 11 раз (3 + 3 + 5) из 399, то есть в 2,8% случаев, то возникает вопрос, действует ли отрицательная часть этой шкалы. Возможно, что попадание в эту часть шкалы – явление чисто случайное. Проверим это предположение.

Будем считать, что если вероятность P попадания в конец шкалы не превышает 5% при достаточно малом уровне значимости ($\alpha = 0,05$ или $\alpha = 0,01$), то наблюдаемые попадания ответов являются случайными и соответствующая часть шкалы «не работает». Для этого границы доверительного интервала,

построенного по имеющейся частоте для вероятности попадания в конец шкалы, сравним со значением 5%. Если значение 5% оказывается выше границ этого интервала, то следует признать, что проверяемая часть шкалы «не работает».

- *Неравномерное использование отдельных пунктов шкалы.*

Случается, особенно при использовании упорядоченных шкал, градации которых сопровождаются словесными описаниями, что некоторое значение переменной (признака) систематически выпадает из поля зрения респондентов, хотя соседние градации, характеризующие более низкую и более высокую степень выраженности признака, имеют существенное наполнение. Так, если конфигурация распределения ответов на вопрос с четырьмя упорядоченными градациями такая, как на рисунке 1.1, то, видимо, шкала неудачно сформулирована.

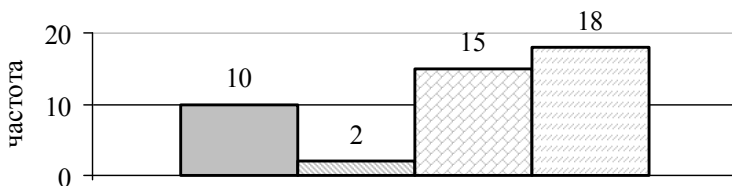


Рис. 1.1. Иллюстрация конфигурации распределения ответов на вопрос с четырьмя упорядоченными градациями

Значительное наполнение двух соседних пунктов («1» и «3») свидетельствует о «захвате» части голосов из плохо сформулированного пункта «2».

Аналогичная картина наблюдается и в том случае, когда респонденту предлагают шкалу, имеющую слишком большую дробность: будучи не в состоянии оперировать всеми градациями шкалы, респондент выбирает лишь несколько базовых. Например, зачастую десятибалльную шкалу респонденты расценивают как некоторую модификацию пятибалльной, предполагая, что «10» соответствует «5», «8» – «4», «5» – «3» и т. д. При этом базовые оценки используются значительно чаще, чем другие.

Для выявления указанных аномалий равномерного распреде-

ления по шкале можно предложить следующее правило: для достаточно большой доверительной вероятности ($1 - \alpha \geq 0,99$) и, следовательно, в достаточно широких границах наполнение каждого значения не должно существенно отличаться от среднего из соседних наполнений.

Соответствующий статистический критерий таков:

$$\chi^2 = \frac{(n_i - \tilde{n}_i)^2 (2n - 1)}{(n_i + \tilde{n}_i)(2n - n_i - \tilde{n}_i)}.$$

Эта величина имеет *Хи-квадрат* распределение с одной степенью свободы ($df = 1$). Здесь i - номер значения признака, который подвергается анализу; n_i - наблюдаемая частота для этого значения.

$$\tilde{n}_i = \frac{n_{i-1} + n_{i+1}}{2} - \text{ожидаемая частота, как средняя из двух соседних.}$$

Например, рассмотрим случай измерения в десятибалльной шкале ряда ценностей типа «любимая работа», «материальный достаток», «здоровье» и т. д. При 45 испытуемых и 14 предложенных ценностях получены 623 оценки, распределение которых выглядит так.

Таблица 1.3

Распределение, демонстрирующее измерение ценностей по десятибалльной шкале

a_i	10	9	8	7	6	5	4	3	2	1
n_i	167	67	90	60	45	81	33	35	28	17

$$n = \sum_{i=1}^{10} n_i = 623.$$

Анализируя таблицу, можно предположить, что значения шкалы 9, 7, 5 не удовлетворяют требованию равномерности. Так, для оценки $a_i = 9$ наблюдаемая частота $n_9 = 67$, ожидаемая

же частота $\tilde{n}_i = \frac{167 + 90}{2} = 128,5$.

Подставим данные значения в формулу χ^2 и получим расчетную величину $\chi^2 = 22,93$. Поскольку $\chi^2 = 22,93 > \chi_{кр}^2 = 6,63$ ($\alpha = 0,01$), то следует признать различие между наблюдаемой и ожидаемой частотами значимым. Следовательно, частота 67 для оценки $a = 9$ слишком мала по сравнению с соседними.

Аналогичные расчеты проводятся для пунктов шкалы $a = 7$ и $a = 5$; частота пункта 7 ($n_7 = 60$) не противоречит выдвинутому требованию равномерности; частота оценки 5 ($n_5 = 81$) слишком велика по сравнению с соседними и, таким образом, противоречит требованию равномерности.

Процедура, направленная на выявление несистематических (случайных) ошибок, – проверка *устойчивости* измерения. О высокой надежности шкалы можно говорить лишь в том случае, если повторные измерения при ее помощи одних и тех же объектов дают сходные результаты. Устойчивость проверяется на одной и той же выборке исследуемых объектов. Линейной мерой несовпадения оценок является *средняя арифметическая ошибка*, показывающая средний сдвиг в ответах в расчете на одну пару последовательных наблюдений. Для этого повторное исследование проводится таким образом, чтобы расстояние по времени одного от другого исследования отличалось настолько, чтобы:

1. Респонденты не смогли воспроизвести ответы.
2. На социальный процесс не повлияли внешние факторы.

Поэтому в зависимости от объекта исследования промежутки между двумя исследованиями выбирают 2–3 недели.

Результаты заносятся в специальную таблицу по типу:

Таблица 1.4

Проба 1	Проба 2			
	X_1	X_2		X_n
X_1	n_{11}	n_{12}		n_{1n}
X_2	n_{21}	n_{22}		n_{2n}
...	...	...	...	...
X_n	n_{n1}	n_{n2}	...	n_{nn}

В качестве меры устойчивости шкалы определяют несколько коэффициентов.

1. Показатель абсолютной устойчивости шкалы

$$W = \frac{\sum_{i=j=1}^k n_{ij}}{n} = \frac{n_{11} + n_{22} + \dots + n_{kk}}{n},$$

где • n_{ii} – количество совпадающих и в первом и во втором опросе;

• n – количество опрошенных.

$W_{\max} = 1$ – в случае, когда различий между первым и вторым опросом нет, то есть все ответы совпадают. Данный коэффициент применяется в основном для качественных признаков номинальной шкалы; для всех же остальных необходимо считать коэффициенты несовпадающих ответов.

Для описания количественных признаков применяются показатели неустойчивости, то есть величины ошибки, учитывающие не только факт несовпадения ответов, но и степень этого несовпадения. Линейной мерой несовпадения оценок является средняя арифметическая ошибка, показывающая средний сдвиг в ответах в расчете на одну пару последовательных наблюдений:

$$|\Delta| = \frac{1}{n} \sum_{i=1}^n |x_i^{II} - x_i^I|,$$

где • x_i^{II} и x_i^I – ответы по анализируемому вопросу i -порядка в I и II пробах соответственно.

Обычно используется вместо этого показателя средняя квадратическая ошибка:

$$S = \sqrt{\frac{1}{2n} \sum_{i=1}^n (x_i^{II} - x_i^I)^2}.$$

В качестве показателя для нормирования абсолютной ошибки можно использовать максимально возможную ошибку в рассматриваемой шкале. Если число делений шкалы k , тогда $\Delta_{\max}$ равно разнице между крайними значениями шкалы ($x_{\max} - x_{\min}$), то есть $k-1$, и относительная ошибка имеет вид:

$$\Delta_{\text{отн}} = \frac{|\Delta|}{\Delta_{\text{max}}},$$

где $|\Delta|$ – средняя арифметическая ошибка измерения.

Хотя исключить случайные погрешности отдельных измерений невозможно, математическая теория случайных явлений позволяет уменьшить влияние этих погрешностей на окончательный результат измерений.

Повышение устойчивости измерения. Для решения этой задачи необходимо выяснить различительные особенности пунктов используемой шкалы, что предполагает четкую фиксацию респондентами отдельных значений: каждая оценка должна быть строго отделена от соседней. На практике это означает, что в последовательных пробах респонденты практически повторяют свои оценки. Следовательно, высокой различимости делений шкалы должна соответствовать малая ошибка.

Эту же задачу можно описать в терминах чувствительности шкалы, которая характеризуется количеством делений, приходящихся на одну и ту же разность в значениях измеряемой величины, то есть чем больше градаций в шкале, тем больше ее чувствительность. Однако чувствительность нельзя повышать простым увеличением дробности, ибо высокая чувствительность при низкой устойчивости является излишней (например, шкала в 100 баллов, а ошибка измерения ± 10 баллов).

Но и при малом числе градаций, при низкой чувствительности, может быть низкая устойчивость, и тогда следует увеличить дробность шкалы. Так бывает, когда респонденту навязывают категорические ответы «да», «нет», а он предпочел бы менее жесткие оценки. И потому он выбирает в повторных испытаниях иногда «да», иногда «нет» для характеристики своего нейтрального положения. Следует найти некоторое оптимальное соотношение между чувствительностью и устойчивостью.

Каждый социолог знает, что ни одно социологическое исследование, основанное на выборочном методе, не может быть лишено случайных ошибок, а каждое эмпирическое наблюдение выражается в численном виде с большей или меньшей долей

ошибки. Распространение результатов выборочного исследования на всю генеральную совокупность невозможно без использования таких понятий, как допустимая «*ошибка выборки*» и «*доверительный интервал*». Первое понятие используется для оценки отклонения выборки от генеральной совокупности. Последнее означает, что существует определенный интервал вокруг числового значения той или иной вариации признака (для выборки), среди которых есть истинное (для генеральной совокупности) значение. Оба эти понятия неразрывно связаны между собой [5, с. 12].

Некоторые факторы могут вызвать одновременно и систематические, и случайные ошибки. Так, включая и выключая секундомер, мы можем создать небольшой нерегулярный разброс моментов пуска и остановки часов относительно движения маятника и внести тем самым случайную ошибку. Но если к тому же мы каждый раз торопимся включить секундомер и несколько запаздываем выключить его, то это приведет к систематической ошибке. Следует иметь в виду, что если случайная погрешность, полученная из данных измерений, окажется значительно меньше погрешности, определяемой точностью прибора, то, очевидно, что нет смысла пытаться еще уменьшить величину случайной погрешности – все равно результаты измерений не станут от этого точнее. Наоборот, если случайная погрешность больше приборной (систематической), то измерение следует провести несколько раз, чтобы уменьшить значение погрешности для данной серии измерений и сделать эту погрешность меньше или одного порядка с погрешностью прибора.

Обоснованность измерения. Наиболее сложный вопрос надежности измерения – его *обоснованность*. Обоснованность шкалы заключается в том, что с ее помощью целенаправленно измеряют вполне определенное свойство или признак, не смешивая его с другими. Обоснованность связана с доказательством того, что измерено вполне определенное заданное свойство объекта, а не какое-либо другое, более или менее не него похожее.

В. А. Ядов, один из известнейших российских социологов, автор многочисленных публикаций, посвященных методологии

и методам анализа социологических данных, в одной из своих работ обращает внимание на один интересный момент: в зарубежной и отечественной (особенно в психологической) литературе вместо термина «обоснованность» часто используется как его аналог понятие «валидность». Однако в английском языке слово *reliability* (обоснованность) подчеркивает возможность полагаться на кого-либо, в данном случае доверять полученной информации благодаря тому, что она адекватна объекту измерения, а *validity* семантически имеет оттенок устойчивости, прочности полученной информации. Поэтому термин «валидность» правильнее было бы соотносить не с обоснованностью, а с устойчивостью данных измерения [7].

Предположим, при опросе телезрителей им предлагают указать, каким из перечисленных в прилагаемом списке передач телевидение уделяет слишком много, достаточно и слишком мало времени. Если с помощью этой трехчленной шкалы исследователь намерен фиксировать среднее время, отводимое телепередачам, его измерение будет необоснованным. В действительности он измеряет отношение людей к данным передачам, а не объем времени, отводимого для их трансляции. Обоснованное измерение объема времени на передачи разного типа – документальный анализ «сетки» программ телевидения [7, с. 205].

Проверка обоснованности шкалы предпринимается лишь после того, как установлены достаточные правильность и устойчивость измерения исходных данных. Проверка обоснованности – достаточно сложный процесс, как правило, не до конца разрешимый. И поэтому нецелесообразно сначала применять трудоемкую технику для выявления обоснованности, а после этого убеждаться в неприемлемости данных вследствие их низкой устойчивости. Обоснованность данных измерения – это доказательство соответствия между тем, что измерено, и тем, что должно было быть измерено. Некоторые исследователи предпочитают исходить из так называемой наличной обоснованности, то есть обоснованности в понятиях использованной процедуры. Рассмотрим формальные подходы к выяснению уровня обоснованности методики. Их можно разделить на три группы:

1. *Конструирование типологии в соответствии с целями исследования на базе нескольких признаков.* Использование контрольных вопросов, которые в совокупности с основными дают большее приближение к содержанию изучаемого свойства, раскрывая различные его стороны.

2. *Использование параллельных данных.* Нередко целесообразно разработать два равноправных приема измерения заданного признака, что позволяет установить обоснованность методов относительно друг друга, то есть повысить общую обоснованность путем сопоставления двух независимых результатов. Классифицируем параллельные процедуры в зависимости от соотношения методов и исполнителей: а) несколько методов – один исполнитель; б) один метод – несколько исполнителей; в) несколько методов – несколько исполнителей.

3. *Судейские процедуры.* Исследователи обращаются к определенной группе людей с просьбой выступить в качестве судей или компетентных лиц. Им предлагают набор признаков, предназначенный для измерения изучаемого явления, и просят оценить правильность отнесения каждого из признаков к этому объекту. Совместная обработка мнений судей позволит присвоить признакам веса или, что то же самое, шкальные оценки в измерении изучаемого явления.

При этом очень важным представляется следующее:

- внимательно проанализировать состав судей с точки зрения адекватности их жизненного опыта и признаков социального статуса соответствующим показателям обследуемой генеральной совокупности;
- выявить эффект индивидуальных отклонений в оценках судей относительно общего распределения оценок;
- следует оценить не только качество, но и объем выборочной совокупности судей.

В завершении данного подраздела подытожим, что в процессе отработки инструментов измерения со стороны их надежности целесообразна следующая последовательность основных этапов работы:

1. Предварительный контроль обоснованности методов

измерения первичных данных на стадии проб методики. Здесь проверяется, насколько информация отвечает своему назначению по существу и каковы пределы последующей интерпретации данных. Для этой цели достаточны небольшие выборки в 10–20 наблюдений с последующей корректировкой структуры методики.

2. Пилотаж методики и тщательная проверка устойчивости исходных данных, в особенности итоговых показателей, индексов, многомерных шкал и т. п. На этом этапе нужна выборка не менее 100 человек, представляющая микромодель реальной совокупности обследуемых с учетом представительства по существенным характеристикам объекта исследования.

3. В период общего пилотажа осуществляются все необходимые операции, относящиеся к проверке уровня обоснованности. Результаты анализа данных генерального пилотажа приводят к усовершенствованию методики, к доработке всех ее деталей и в итоге – к получению окончательного варианта методики для основного исследования.

4. В начале основного исследования желательно провести проверку используемого варианта методики на устойчивость с тем, чтобы рассчитать точные показатели ее устойчивости. Последующее уточнение границ обоснованности проходит через весь анализ самого исследования.

Вопросы и задания для самоконтроля

1. *Дайте определение термину «измерение». Что и зачем измеряется в точных и естественных науках?*

2. *Что общего и особенного в таких терминах, как «признак» и «переменная»?*

3. *Что, зачем и с помощью каких средств измеряют социологи?*

4. *Что такое шкала? Какие бывают шкалы, приведите примеры.*

5. *Проанализируйте возможности и специфические особенности шкал, применяемых в социологических исследованиях.*

6. *Какие типы шкал используются социологами? От чего зависит тип шкалы?*

7. Может ли один и тот же признак быть измерен с помощью разных типов шкал? Приведите пример.

8. Какой смысл закладывается в термин «надежность» измерения в социологии?

9. Составные части надежности измерения.

10. Тождественны ли термины «ошибка» и «погрешность» для социолога? Дайте развернутый ответ на вопрос, обоснуйте его.

11. Назовите и опишите основные типы ошибок в социологическом исследовании. Каковы их основные отличия?

12. Предложите «социологическую» трактовку термина «правильность» измерения.

13. Определите, что такое устойчивость измерения?

14. С разрешением каких ошибок связан вопрос об устойчивости измерения?

15. В чем заключается суть обоснованности измерения? С помощью каких методов подтверждается обоснованность измерения?

16. Проверьте на устойчивость (с применением соответствующей формулы) осуществленное измерение, результаты которого зафиксированы в таблице, представленной ниже.

Проба I	Проба II					Сумма
	1	2	3	4	5	
1	3	5	1			9
2		3	1	1		5
3		7	6	2	2	17
4	1	3	4	6	1	15
5		1		1	2	4
Σ	4	19	12	10	5	50

Раздел II. АГРЕГИРОВАНИЕ СОЦИОЛОГИЧЕСКИХ ДАННЫХ И СООТВЕТСТВУЮЩИЕ МАТЕМАТИЧЕСКИЕ ПРОЦЕДУРЫ

2.1. Частотные распределения

Агрегирование – укрупнение тех или иных показателей посредством их объединения в единую группу, то есть сведение частных показателей к обобщенным. В результате агрегирования социологических данных исследователь получает важные «синтетические измерители», объединяющие в себе множество частных показателей. Агрегирование осуществляется посредством суммирования, группировки или других способов сведения частных показателей в обобщенные. Измеряя характеристики объекта, исследователь собирает первичный статистический материал, который в процессе агрегирования систематизируется и обобщается для выявления характерных черт, свойств, типов социальных явлений, для обнаружения закономерностей изучаемых процессов и проверки исследовательских гипотез. В основе используемых методов обработки первичных данных, их агрегирования, упорядочения лежит, главным образом, статистическая группировка, а также составление статистических таблиц [1, с. 32].

Как говорилось в предыдущем разделе, проводя исследование, социолог изучает признаки, которые характеризуют исследуемую совокупность. На основе анализа этих признаков исследователь делает выводы относительно того или иного социального явления, процесса и т. п.

Признак – случайная величина, которая варьируется, то есть имеет определенное число вариаций, принимает различные значения. Для каждого значения этой величины, по прохождении полевого этапа исследования, будет известна частота встречае-

мости. Другими словами социологу будет известно одномерное распределение вероятностей, которые задают эту случайную величину.

Результаты измерения, как правило, изначально имеют произвольный и хаотичный порядок. В такой форме полученные данные неудобны для анализа и выявления закономерностей. Первичная обработка статистических данных состоит в упорядочении данных (по возрастанию или убыванию), подсчете некоторых показателей, характеризующих эти значения, в группировании данных.

Ряды распределения – это ряды абсолютных и относительных чисел, которые характеризуют распределение единиц совокупности по качественному (атрибутивному) или количественному признаку. *Зарегистрированные в результате наблюдения индивидуальные значения изучаемого варьирующего признака образуют так называемый **первичный ряд**.*

Ряд значений признака, или вариант, полученных вследствие массового обследования однородных вещей или явлений, размещенных в порядке возрастания или убывания их величин, вместе с соответствующими частотами (или относительными частотами) называют **вариационным рядом**.

Если в вариационном ряде значения признака (варианты) заданы в виде отдельных конкретных чисел, то такой ряд называют **дискретным**.

Если в вариационном ряде значения признака заданы в виде интервалов, то такой ряд называют **интервальным**.

Еще есть такое понятие, как «**динамический** (или временной) ряд», показывающее движение признака (изменение его частот) во времени, то есть изменение его в связи с переходом от одного момента или периода времени к следующему. Изучение динамических рядов позволяет установить закономерность в развитии данного явления или признака, определить складывающиеся тенденции и выявить различные колебания, вариации, отклонения.

Число случаев, в которых встречается то или иное значение признака (варианта), называют **абсолютной частотой** этого

значения. Результаты социологического исследования фиксируются как статистические наблюдения, которые регистрируются прежде всего в форме первичных абсолютных величин или абсолютных частот.

Например, если мы изучаем распределение респондентов по профессиональному признаку (качественному), для конкретной выборки, объем которой равен 80 человек ($n = 80$), то ряд распределения абсолютных частот может выглядеть следующим образом:

1. Учитель – 10 человек.
2. Врач – 11 человек.
3. Повар – 9 человек.
4. Водитель – 20 человек.
5. Продавец – 30 человек.

Абсолютные частоты в качестве единицы измерения всегда имеют «штук», «человек», «часов», «килограмм» и т. п.

Но более удобным для восприятия и дальнейшего анализа (например, сравнительного), является представление данных в виде относительных частот. **Относительная частота** – доля или процент объектов, обладающих данным значением признака, по отношению к объему выборки.

Абсолютные частоты, с помощью элементарных математических вычислений из школьной программы, легко переводятся в относительные. В случае с примером, приведенным выше, можем иметь следующую картину:

1. Учитель – 12,5 %.
2. Врач – 13,75%.
3. Повар – 11,25%.
4. Водитель – 25%.
5. Продавец – 37,5%.

Если значения признаков выражены в относительных числах, то эти значения именуются **частостями**. Распределение по признаку профессиональной принадлежности – это распределение по качественному признаку. Однако эта же совокупность может быть распределена по количественному признаку, например, по возрасту. Диапазон возможных вариантов такой

переменной, как возраст, очень широк. В предыдущем разделе мы говорили о том, что эту переменную можно представить в виде метрической (числовой) шкалы, а градации еще и выстроить по порядку. Такое представление нам может быть необходимо для того, чтобы расширить спектр возможных математических операций с этой шкалой¹⁰. Однако вспомним, что представленная таким образом переменная возраста – это непрерывная количественная переменная. Она может принимать огромное множество различных значений. Чем больше число этих значений, тем сложнее человеческому мозгу справиться с их обработкой. Например, мы с легкостью можем представить трех- или четырехугольник. Однако «двенадцатиугольник» уже представить довольно сложно, а стоугольник – вообще невозможно. Хотя каждый из нас может вполне успешно описать эти фигуры, дать им определение. А для того чтобы мы все-таки смогли представить, вообразить объект исследования, необходимо воспользоваться теми или иными методами агрегирования данных, то есть их «укрупнения» посредством объединения по группам. Агрегирование, как правило, осуществляется посредством группировки, что позволяет информацию об объекте исследования представить в «компактной» форме, такой, которая легко поддавалась бы нашему восприятию и воображению. Для этого осуществляется группировка данных.

Группировка – *разбиение* совокупности на группы, однородные по какому-либо признаку (а с точки зрения отдельных единиц совокупности, наоборот, – *объединение* отдельных единиц в однородные группы). Метод группировки основывается на следующих категориях: группировочный признак, интервал группировки и число групп.

Группировочный признак – это признак, по которому происходит объединение отдельных единиц совокупности в однородные группы (в нашем случае группировочным признаком выступает признак возраста).

¹⁰ Вспомним, что номинальная шкала – самая ограниченная в плане допустимых операций.

Интервал очерчивает количественные границы групп. Как правило, он представляет собой промежуток между максимальными и минимальными значениями признака в группе. Интервалы бывают:

- *равные*, когда разность между максимальным и минимальным значениями в каждом из интервалов одинакова;
- *неравные*, когда, например, ширина интервала постепенно увеличивается, а верхний интервал часто не закрывается вовсе;
- *открытые*, когда имеется только либо верхняя, либо нижняя граница;
- *закрытые*, когда имеются и нижняя, и верхняя границы.

Определение числа групп (числа интервалов) и их границ подчинено ряду условий:

А. Число групп (***k***) детерминируется уровнем однородности/неоднородности группировочного признака. Чем сильнее неоднородность, чем больше варьирует признак, тем больше должно быть число групп. Например, если мы опросили только школьников 10–11 классов, диапазон изменения значений признака будет небольшим, и, скорее всего, этот диапазон мы разобьем на два возрастных интервала, либо вообще не будем разбивать (все зависит от целей исследования).

Выбор числа интервалов группировки можно осуществить при помощи таблицы 2.1, приведенной ниже:

Таблица 2.1

Таблица определения числа интервалов

Объем выборки, <i>n</i>	Число интервалов, <i>k</i>
(25–40)	(5–6)
[40–60)	[6–8)
[60–100)	[8–10)
[100–200)	[10–12)
Больше 200	[12–15)

Б. Число групп должно отражать реальную структуру изучаемой совокупности.

В. Не допускается выделение пустых групп. Если проблема пустых групп все же возникает, при проведении структурных группировок используют неравные интервалы.

Г. Необходимо четко обозначить границы каждого интервала. Для начала нужно найти ширину каждого из интервалов. Шириной интервала (h) называют разность между верхней и нижней границами интервала.

В случае равных интервалов их ширина находится по следующей формуле:

$$h = \frac{x_{\text{макс}} - x_{\text{мин}}}{k - 1},$$

где • h – ширина интервалов;

• $x_{\text{макс}}$ и $x_{\text{мин}}$ – максимальная и минимальная варианты выборки ($x_{\text{макс}}$ и $x_{\text{мин}}$ находятся непосредственно по таблице исходных данных);

• k – число интервалов.

Прибавив к этой величине ширину интервала, найдем нижнюю границу второго интервала (x_{n2}):

$$x_{n2} = x_{n1+h}.$$

Это будет одновременно и верхняя граница предыдущего (первого) интервала. В этой связи возникает резонный вопрос: «К какому интервалу относить варианту, находящуюся на стыке двух интервалов?». Однозначного ответа на этот вопрос нет. Авторы многих учебных пособий подчеркивают, что такие варианты могут быть с одинаковыми основаниями отнесены к любому из соседних интервалов, и исследователь сам решает, как поступить в этой ситуации. Тем не менее, мы склонны придерживаться правила, изложенного в следующем абзаце, которое для многих является негласным.

Д. Если в интервальном вариационном ряде в двух последовательных интервалах верхнее предельное значение признака одного интервала равняется нижнему предельному значению второго, условно будем считать, что это число принадлежит *второму интервалу*. В письменной форме это выражается

с помощью широко известных условных обозначений – «круглой» и «квадратной» скобки, т. е. $[x_{n2}; x_{n1+h})$.

Рассмотрим еще два символа, которые нам могут пригодиться при записи интервалов: $-\infty$ (минус бесконечность) и $+\infty$ (плюс бесконечность). Они не являются числами и вводятся лишь для удобства записи, когда нам неизвестны, либо мы сознательно не хотим задавать самую нижнюю и самую высокую границы всего диапазона значений.

Проиллюстрируем все сказанное примером. Всех респондентов, которых в начале данного параграфа мы распределили по профессиональному признаку, можно распределить и по признаку возрастному. Признак возраста зададим метрической шкалой. При этом представим, что диапазон распределения всех возможных значений признака довольно широк: мы опрашивали все население трудового возраста, начиная с тех, кто только что окончил школу, заканчивая теми, кто в ближайший год собирается выходить на пенсию. Всего нами было опрошено 80 человек ($n = 80$). Исходя из таблички определения интервалов, мы можем разбить весь диапазон на 8–10 интервалов. Самому младшему респонденту, в идеале, должно быть 17 лет, самому старшему – 55. Следовательно, крайними (нижним и верхним) значениями диапазона распределения должны быть, соответственно 17 и 55. Но ведь известно, что некоторые молодые люди оканчивают школу в более раннем возрасте, например, в 15 или 16 лет. Такие же неточности могут быть и в отношении предпенсионного возраста. С учетом этого мы сознательно не ограничиваем крайнюю нижнюю границу первого интервала и крайнюю верхнюю границу последнего. Того человека, которому полных 22 года (27, 32, 37 и т. д.), относим не к первому, а ко второму интервалу. В итоге имеем:

- 1) моложе 22 лет – 5 человек (6,25%),
- 2) 22–27 года – 7 человек (8,75%),
- 3) 27–32 лет – 10 человек (12,5%),
- 4) 32–37 лет – 11 человек (13,75%),
- 5) 37–42 лет – 18 человек (22,5%),

- 6) 42–47 лет – 12 человек (15%),
 7) 47–52 лет – 9 человек (11,25%),
 8) 52 и старше – 8 человек (10%).

Для удобства занесем все эти цифры в таблицу 2.2, представленную ниже.

Таблица 2.2

Распределение респондентов по возрасту
 (абсолютные и относительные для $n = 80$)

x_{ni} ; x_{ni+h} Возраст	$(-\infty; 22)$	$[22; 27)$	$[27; 32)$	$[32; 37)$	$[37; 42)$	$[42; 47)$	$[47; 52)$	$[52; +\infty)$
Абсолютная частота n_i (человек)	5	7	10	11	18	12	9	8
Относительная частота, частость v_i (в % ко всем опрош.)	6,25	8,75	12,5	13,75	22,5	15	11,25	10

Из таблицы мы можем увидеть, насколько часто повторяется та или иная варианта, какова частота попадания респондентов в тот или иной интервал. Очевидно, что, например, самая большая частота у интервала $[37; 42)$ и равна она 18 человек или 22,5%.

Помимо этого, рассмотрим так называемые **кумулятивные** вариационные ряды. В таких рядах вместо частот или относительных частот определенных вариантов (или интервалов) записаны **накопленные** (кумулятивные) **частоты** или относительные накопленные частоты. Накопленная частота – число, полученное последовательным суммированием частот в направлении от первого интервала к последнему, до того интервала включительно, для которого определяется накопленная частота. Вычисление накопленных частот необходимо для определения некоторых характеристик положения, например, медианы и других квантилей. Более подробно об этом будет сказано в параграфе **2.3** данного учебника. Просчитаем накопленные частоты для нашего примера с возрастом респондентов. Результаты представлены в таблице 2.3.

Таблица 2.3

Распределение респондентов по возрасту
(абсолютные, относительные, кумулятивные частоты для $n = 80$)

x_{ni} ; x_{ni+h} Возраст	$(-\infty; 22)$	$[22; 27)$	$[27; 32)$	$[32; 37)$	$[37; 42)$	$[42; 47)$	$[47; 52)$	$[52; +\infty)$
Абсолютная частота n_i (человек)	5	7	10	11	18	12	9	8
Относительная частота, частость V_i (% ко всем опрош.)	6,25	8,75	12,5	13,75	22,5	15	11,25	10
Кумулятивная частота n_i^H (человек)	5	12	22	33	51	63	72	80
Кумулятивная частость V_i^H (% ко всем опрош.)	6,25	15	27,4	41,15	63,65	78,65	89,9	100

В заключение отметим, что группировка данных, так сказать, предварительная процедура их агрегирования, она формирует основу для последующей сводки и анализа данных. Конечной целью агрегирования является представление множественных частных значений и показателей в одном общем показателе. Агрегированные показатели – обобщенные, синтетические «измерители». В социологии, как правило, используется три основных метода агрегирования данных, заимствованных из описательной статистики: 1) табличное представление; 2) графическое изображение; 3) расчет статистических показателей. К более подробному рассмотрению этих методов мы предлагаем обратиться в последующих параграфах данного раздела.

2.2. Перекрестная классификация и табличное представление социологической информации

В предыдущем параграфе речь шла о группировке как методе обработки социологических данных, позволяющем обеспечить первичное обобщение данных, представление их в более упорядоченном виде. Однако есть еще один, не менее важный метод обработки данных – метод классификации.

Классификация – это систематическое распределение явлений и объектов по определенным группам, классам, разрядам на основании их сходства и различия [16].

Социологическое исследование редко ограничивается представлением одной переменной в виде ряда интервалов и категорий. Как правило, исследователь заинтересован в обнаружении связи двух и более переменных. Для данной цели более эффективна классификация по нескольким переменным – двухвариантная или поливариантная классификация.

Статистический метод располагает широким разнообразием приемов разной степени сложности для анализа подобных отношений, и одним из простейших приемов является метод перекрестной классификации или «матричный» метод. Согласно этому методу, ряд объектов группируется не по одному признаку, а по двум и более одновременно. Полученные в результате частоты названы совместными, потому что показывают число совместных событий, дают ключи к пониманию того, каким образом можно связать эти переменные, и могут дать больший эффект, чем более квалифицированная, сложная корреляционная техника.

Типы перекрестной классификации и их интерпретация.

В таблице 2.4 для 28 объектов, в качестве которых выступают области Украины, Автономная Республика Крым, отдельно г. Севастополь и г. Киев и Украина в целом, построена перекрестная классификация, или «матрица», в которой представлено распределение численности докторов и кандидатов наук, в зависимости от их географического местоположения на карте Украины.

Таблица 2.4

**Численность преподавательского состава
на 100 студентов дневной формы обучения
(распределение по регионам, в %)**

<i>Область</i>	<i>Доктора наук</i>	<i>Кандидаты наук</i>
Автономная республика Крым	0,9	4,5
Севастополь	0,9	4,2
Винницкая	0,8	5,3
Волинская	0,4	4,3
Днепропетровская	0,8	4,1
Донецкая	0,6	3,5
Житомирская	0,4	3,6
Закарпатская	1,1	4,2
Запорожская	0,6	4,5
Ивано-Франковская	1	5
Киев	1,2	5,5
Киевская	0,5	2,7
Кировоградская	0,3	3
Луганская	0,6	4,2
Львовская	0,9	5,5
Николаевская	0,3	3,1
Одесская	1,1	5,5
Полтавская	0,5	4,2
Ровенская	0,4	4,7
Сумская	0,6	4,2
Тернопольская	0,6	3,3
Харьковская	0,9	5,4
Херсонская	0,5	3,3
Хмельницкая	0,5	4,4
Черкасская	0,5	3,8
Черниговская	0,4	3,7
Черновицкая	1	5,8
Украина – всего	0,8	4,6

Данный пример служит наглядной иллюстрацией того, как посредством метода классификации (перекрестной) мы можем осуществить сравнительный анализ распределения преподавательских кадров на востоке, западе, юге, севере и в центре страны.

Таблица 2.4 построена по качественной переменной и двум количественным. Однако аналогичным образом можно построить таблицу по двум качественным признакам (см. табл. 2.5).

Таблица 2.5

Распределение ответов респондентов в зависимости от их пола на вопросы анкеты: «На что Вы в наибольшей мере полагаетесь в своей жизни?»
(в % к опрошенным)¹¹

<i>Категории</i>	<i>Мужчины</i>	<i>Женщины</i>
На личные знания, способности и силы	72	61
На поддержку со стороны собственной семьи	20	27
На бога	16	25
На удачу, на судьбу	17	13
На поддержку со стороны родителей (родительской семьи)	13	13
На поддержку со стороны государства	5	6
На поддержку со стороны друзей и знакомых	6	5
Другое	1	1
Трудно ответить	1	2

Как видим из примера, приведенного выше, нет никакой возможности оценить значение цифр, представленных в таблице, осуществить полноценный сравнительный анализ распределений ответов, в зависимости от половой принадлежности. Измерение осуществлялось по шкале с совместимыми альтернативами, следовательно, сумма всех ответов превышает 100%. Поэтому необходимо, чтобы для каждого ряда процентов было указано исходное количество случаев – N с тем, чтобы в распределении нашли свое отражение действительные отношения.

¹¹ Моніторинг громадської думки населення України. Інформ. бюл. / Укр. ін-т соц. дослідж., Центр «Соц. Моніторинг». – К., 1999. – № 1. – 40 с.

Может быть и такое, что процентные отношения, расположенные в ячейках матрицы (пересечениях строк и столбцов), вычисляются исходя не из полной суммы случаев, а лишь из случаев, относящихся к данной строке или столбцу («маргинальных сумм»). Обычно маргинальные процентные отношения выделяют какую-либо независимую переменную.

Как можно было заметить, результаты группировки и классификации данных социологического исследования довольно часто оформляются в виде статистических таблиц, где излагаются в наглядно-рациональной форме. Однако, следует подчеркнуть, что не всякая таблица может быть названа статистической. Табличные формы календарей, тестовых и опросных листов, таблица умножения не являются статистическими [2, с. 33; 8].

Статистическая таблица – это цифровое выражение итоговой характеристики всей наблюдаемой совокупности или ее составных частей по одному или нескольким существенным признакам. Статистическая таблица содержит два элемента: подлежащее и сказуемое [2, с. 35; 9]. Каждую статистическую таблицу можно рассматривать как «оконченную мысль» исследователя. И каждая из них имеет свое подлежащее и сказуемое.

Подлежащее статистической таблицы – перечень групп или единиц, составляющих исследуемую совокупность единиц наблюдения.

Сказуемое статистической таблицы – цифровые показатели, с помощью которых дается характеристика выделенных в подлежащем групп и единиц.

Различают *простые, групповые и комбинационные таблицы*. В **простых таблицах**, как правило, содержится справочный материал, где дается перечень групп или единиц, составляющих объект изучения. При этом части подлежащего не являются группами одинакового качества, отсутствует систематизация изучаемых единиц. Сказуемое этих таблиц содержит абсолютные величины, отражающие объемы изучаемых процессов. Групповые и комбинационные таблицы предназначены для научных целей, где, в отличие от простых таблиц, сказуемое составляют средние и относительные величины на основе абсолютных величин.

Групповая таблица – это таблица, где статистическая совокупность разбивается на отдельные группы по какому-либо одному существенному признаку, при этом каждая группа характеризуется рядом показателей. Примером такой группировки может быть разделение семей на группы по месту проживания (сельское и городское), где образуются подгруппы семей по количеству детей. Анализ этих группировок по материалам переписи 1989 года позволил сделать вывод, что большинство семей, независимо от принадлежности к городскому или сельскому населению, имеют только по одному ребенку.

Комбинационная таблица – это таблица, где подлежащее представляет собой группировку единиц совокупности по двум и более признакам, которые распределяются на группы сначала по одному признаку, а затем на подгруппы по другому признаку внутри каждой из уже выделенных групп. Комбинационная таблица устанавливает существенную связь между факторами группировки. Примером комбинационной группировки может быть распределение полиграфических предприятий по трем существенным признакам: степени оснащенности современным полиграфическим оборудованием, степени применения современных технологий и уровню производительности труда. Такого рода статистические таблицы позволяют осуществить всесторонний анализ, но они менее наглядны.

Независимо от количества переменных, для которых строятся таблицы, есть единый перечень общих принципов, которыми при этом следует руководствоваться: количество и величину интервалов необходимо устанавливать таким образом, чтобы это не нарушало и не затуманивало взаимоотношение этих двух переменных. Что касается качественных показателей, признаки должны быть исчерпывающими и взаимоисключающими. Если эти простые принципы будут серьезно приняты во внимание, подготовка таблицы с перекрестной группировкой не составит труда.

В заключении подчеркнем, что при составлении таблиц необходимо соблюдать **общие правила**:

- таблица должна быть легко обозримой;

- общий заголовок должен кратко выражать основное содержание;
- наличие строк «общих итогов»;
- наличие нумерации строк, которые заполняются данными;
- соблюдение правила округления чисел.

2.3. Графическое представление социологической информации

Назначение графика. Весь табличный материал можно представить в графической форме, которая (обычно) более наглядно, чем таблица, выражает картину общего распределения. Таблица частот позволяет с большей легкостью интерпретировать результаты. Главное назначение графика – дать наиболее точное представление о форме частотного распределения – представление, понятное даже совершенно неквалифицированному читателю.

Существует множество видов графического представления, каждое из которых полезно в своем конкретном приложении [4, с. 20–27; 5]. Некоторые из них весьма сложны, но здесь рассматриваются только основные и простейшие: 1) гистограмма; 2) полигон распределения; 3) кумулята; 4) диаграмма полос; 5) статистическая карта; 6) временная диаграмма.

Первые три из указанных шести типов применимы только к количественным данным. Диаграмма полос предназначена, главным образом, для качественных данных; статистическая карта представляет распределение событий по географической площади, а временная диаграмма является графическим вариантом динамического ряда.

Так как графики эквивалентны таблицам, то они должны иметь аналогичные названия и обозначения и подчиняться тем же критериям доступности, простоты и ясности. График нельзя построить до тех пор, пока не будет приготовлена соответствующая таблица.

Построение гистограммы. Гистограмма состоит из ряда соприкасающихся столбцов, высота которых пропорциональна

частоте соответствующего класса событий, а ширина пропорциональна величине интервала группировки переменной. Она является не только графической записью абсолютных частот группировок, но и наглядным изображением значения каждой частоты относительно всех других. В качестве примера предлагаем рассмотреть гистограмму, изображающую частотное распределение сети высших учебных заведений III–IV уровня аккредитации Украины (начиная с 1992 г.), которая представлена на рисунке 2.1. Если строить гистограмму вручную, то делается это на обыкновенной бумаге, линованной «в клетку». В качестве первого шага на сетке линий проводятся под прямыми углами две оси приблизительно равной длины, которые пересекаются вблизи нижнего угла с левой стороны страницы. Границы группировок наносятся на горизонтальной *оси X* (абсцисс), а частоты группировок наносятся на вертикальной *оси Y* (ординат). Однако, прежде чем располагать интервалы группировок на оси абсцисс, необходимо установить, какое число линейных единиц можно поставить в соответствие каждому интервалу группировки. Для улучшения эстетического вида графика и удобства чтения полезно оставлять свободные отрезки в начале и в конце горизонтальной оси. Учитывая это, подсчитывают полное число единиц длины оси и делят его на число интервалов группировок, которые необходимо нанести на график. В результате получается число линейных единиц, предназначенных для каждого интервала. Теперь, начиная от точки пересечения двух осей, можно откладывать интервалы группировок и помечать их границы острыми вертикальными линиями или толстыми штрихами, размещенными соответственно с внешней стороны оси, так как они могут совпасть со следами первоначальной сетки линий. В связи с тем что эти пометки представляют собой общие точки соприкасающихся интервалов, они же обозначают истинные пределы группировок. Эти истинные пределы неизбежно будут дробными, когда пределы округляются до ближайшего целого.

Аналогичным образом устанавливается шкала частот на вертикальной оси. Наибольшую частоту группировки, которую необходимо разместить на графике, делят на число линейных

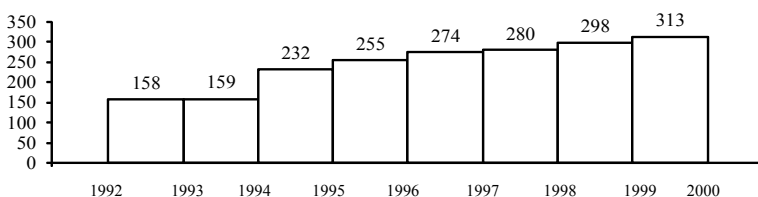


Рис. 2.1. Гистограмма динамики численности
высших учебных заведений III–IV уровня аккредитации Украины
(в абсолютных частотах)

единиц, имеющих на предварительно проведенной оси. Тем самым определяется число событий, соответствующих каждой единице вертикальной оси.

Начиная с нулевой точки – точки пересечения осей, обозначения размещаются через равные интервалы и даются такими величинами (5, 10, 15 и т. д.), которые могут быть ясны и понятны.

Построив систему обозначений обеих осей, можно переходить к построению столбцов гистограммы. Вертикальные границы столбцов проводятся из точек истинных границ, а их высоты определяются частотами соответствующих интервалов.

Как правило, отсчет шкалы частот начинается с нуля, в противном случае не соблюдается необходимая пропорциональность площадей столбцов. Практика показывает, что желательно для лучшего восприятия пользоваться осями, равными по длине. Удлиняя ось абсцисс и укорачивая ось ординат, можно сделать гистограмму длинной и плоской, создавая тем самым впечатление большей вариации переменной. С другой стороны, удлинение оси ординат и укорочение оси абсцисс создает высокие узкие фигуры, которые представляют видимость незначительной вариации.

Неравные интервалы группировки. В случае неравных интервалов необходимо добиться возможности сравнения частот и сохранения их пространственных отношений. Поэтому необходимо разбить промежуток, который шире остальных, на два интервала равной ширины; соответственно этому необходимо

разделить и его частоту на две равные частоты. Затем эти преобразованные подчастоты вычерчиваются на графике. Если бы начертили непреобразованную частоту, то столбец, соответствующий нестандартному интервалу, имел бы значение в два раза большее по площади, чем он имеет на самом деле, и тем самым создавал бы впечатление, не соответствующее действительности.

Обычно, прежде чем строить график по таблице, в которой существуют неравные интервалы группировок, необходимо представить все большие интервалы в виде кратного числа меньших интервалов и на это число разделить соответствующие частоты. Последний шаг приводит к преобразованным частотам, которые затем и вычерчиваются. Эта процедура находится в согласии с тем принципом, что каждое событие в распределении частот представлено равной площадью графика, так что относительные частоты пропорциональны площадям.

Теперь можно оценить практическую полезность правила о том, что интервалы группировки должны быть равной ширины или, в крайнем случае, должны быть составлены из целого числа меньших интервалов. Поэтому невозможно представить таблицу с открытыми интервалами в виде гистограммы, если не прибегать к произвольному ограничению интервалов – что иногда и приходится делать. Отличительной чертой гистограммы является ее схематическая простота. Столбцы более выразительны, чем числа. Они ясно раскрывают относительную плотность событий в каждом интервале и показывают контур распределения.

Природа и построение полигона распределения. Когда события сосредоточены в относительно небольшом количестве широких интервалов, частоты имеют тенденцию резко, скачком обрываться на границе каждого интервала. Разумнее, однако, предположить, что последовательность частот группировок была бы значительно глаже, если бы применяли большее число относительно небольших интервалов. Полигон распределения предназначен дать такое приближение в виде сглаженной кривой, которая, возможно, возникла бы, если бы размеры интервалов стремились к нулю, а число наблюдений неограниченно возрастало [4, с. 20–25; 5].

Полигон распределения можно получить из гистограммы, проводя прямые линии через средние точки верхних частей смежных столбцов. Такое преобразование изображено на рисунке 2.2.

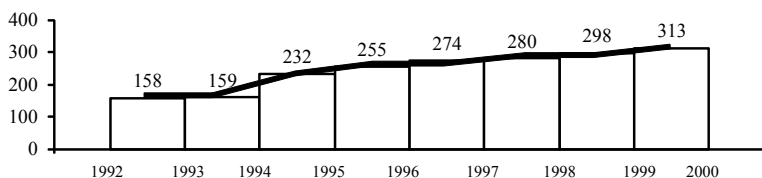


Рис. 2.2. Полигон и гистограмма динамики сети высших учебных заведений III–IV уровня аккредитации Украины (в абсолютных частотах)

На практике строится только один из графиков в зависимости от того, придается ли особое значение «затабулированным» частотам или же тем гипотетическим «точечным» частотам, которые дает полигон распределения. Очевидно, что процедура построения полигона распределения во всем соответствует процедуре построения гистограммы, за исключением конечной стадии. Сначала на линованной бумаге проводятся две оси, затем на базовой линии (оси абсцисс) наносятся интервалы, а шкала частот откладывается вдоль вертикальной оси. Однако на следующей стадии процедуры расходятся. Вместо столбцов наносятся линии, соединяющие точки, расположенные над серединами интервалов на высотах, соответствующих частотам данных интервалов; затем эти точки соединяются отрезками, образуя многоугольник. Нет никакой необходимости продолжать этот многоугольник (полигон) до базовой линии, за пределами области действительных наблюдений. Так, в V-образном распределении, где концентрация событий увеличивается на обоих концах, это было бы даже нелогично. По-видимому, более целесообразно оставить многоугольник открытым на обоих концах и тем самым избежать опасности ложного изображения частот, из которых ни одна фактически не наблюдалась. Иногда

все же график доводят до базовой линии, опуская отрезки в средние точки свободных интервалов на обоих концах. Это объясняется стремлением определить площадь, принадлежащую гистограмме, которая графически изображает сумму частот. Хотя это и достаточно разумно, необходимо все же учитывать, что многоугольник неизбежно нарушает (как правило, преуменьшает) отношение между площадями столбцов и частотами группировок. В связи с этим необходимо помнить, что не существует идеальных статистических приемов, всесторонне и точно представляющих данные. Полигон распределения претендует только на то, чтобы быстро и экономично получить первое приближение теоретической кривой распределения частот величин, сгруппированных в возможно меньшие интервалы. Для целей сравнения можно налагать друг на друга полигоны распределения, представляющие различные распределения одних и тех же переменных, не нарушая их относительных очертаний.

Кумулята. Построение кумуляты. Кумулятивная таблица частот (см. *Табл. 2.6*) дает процент случаев ниже (или выше) каждой данной границы группировки, но не аккумулируют частоты в пределах различных границ, лежащих в пределах вариации переменной. Поэтому она не может давать величины, которые соответствуют всем возможным кумулятивным процентам. Тем не менее часто требуется подобная информация. Так, например, может потребоваться подсчитать процент первоклассников, поступивших в школы г. Харькова с 1998 по 2000 год, или найти такое граничное количество, когда 50% первоклассников поступило в школы в период с 1998 по 2001 год, хотя этот вид информации можно получить из совокупной таблицы с помощью арифметической интерполяции. Эта же информация получается с меньшим усилием и достаточно точно из кумулятивного полигона распределения или кумуляты, как его иногда называют.

Построение кумуляты начинается с вычерчивания осей, как и для простого полигона распределения. Интервалы группировок наносятся на ось абсцисс, а частоты – на ось ординат. Так как все частоты теперь кумулятивные, то верхнее деление вертикальной шкалы будет соответствовать полной сумме частот.

Таблица 2.6

**Прием в первые классы общеобразовательных школ
г. Харькова¹²**

Учебный год	Количество первоклассников: <u>абсолютные</u> частоты (n_i)	Кумулятивные частоты <u>абсолютные</u> (n_i^H)	Количество первоклассников: <u>относительные</u> частоты (v_i в %)	Кумулятивные частоты <u>относительные</u> (v_i^H , %)
1998/1999	14452	14452	32,64%	32,64%
1999/2000	13601	=14452+13601=28053	30,72%	=32,64%+30,72%=63,36%
2000/2001	16221	44274	36,64%	100,00%
Итого	44274		100,00%	

Для быстрого понимания и легкости сравнения кумулятивные частоты обычно выражаются в процентах, так что шкала частот простирается от 0 до 100. Соответственно двум типам таблиц кумулятивных частот: «меньше чем» и «или более» – существует два типа кумулят. Но эти два графика дают идентичную информацию, так что для всех практических целей строиться будет только один. При построении кумуляты «меньше чем» (см. Рис. 2.3) кумулятивные частоты наносятся над истинными верхними пределами интервалов. Эта процедура отличается от процедуры построения простого полигона распределения, где фиксируются средние точки интервалов.

Некоторые авторы пользуются терминологией «меньше чем» и «больше чем», например, «менее чем 8» и «более чем 3». Однако эти термины создают трудности при применении к дискретным данным, если только они не трактуются как непрерывные. Если 33% всех семей состоят менее чем из трех человек, а 45% – более чем из трех, то семья, величиной в три человека, остается неучтенной. Однако когда данные являются непрерывными, то этих трудностей не возникает, так как точка разрыва не занимает

¹² Показники роботи закладів освіти та наукових установ області за 2000 рік // Стат. зб. : за заг. ред. О. Л. Сидоренко, Л. О. Белової, А. С. Доценка. – Х., 2001. – 87 с.

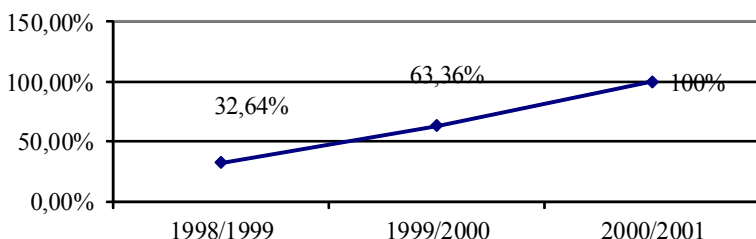


Рис. 2.3. Кумулята (график кумулятивных частот) приема в первые классы общеобразовательных школ г. Харькова

никакой части континуума. Обозначение, принятое в этом тексте, является достаточно гибким для того, чтобы удовлетворить требованиям дискретных данных и не сделать насилия над непрерывными данными.

Благодаря такому способу построения шкал, кумулята «менее чем» будет начинаться в нижнем левом углу и проходить по диагонали в верхний правый угол, в то время как кумулята «или более» начинается в верхнем левом углу и движется диагонально к правому нижнему углу. Когда один или оба конца таблицы являются открытыми, кумулята не будет уже распространяться на всю частотную шкалу и, таким образом, будет казаться неполной. В этом случае можно ограничить кумуляту некоторой удобной точкой, не искажая данных.

Кумулята позволяет расчленять частотное распределение в любых точках в зависимости от необходимости, например, мы с легкостью сможем определить медиану (Me), найдя точку, четко разделяющую все события исследуемой совокупности на две равные части. Или точку, которая разделяет нижние 25% и верхние 75% событий, так называемый, первый квартиль (Q_1), как на рисунке 2.4, представленном ниже.

От отметки в 25% на оси частот мы провели линию, параллельную базовой линии, до тех пор, пока она не пересечет кумуляту «меньше чем». Из этого пересечения опускаем на базовую линию перпендикуляр, который фиксирует первый квартиль

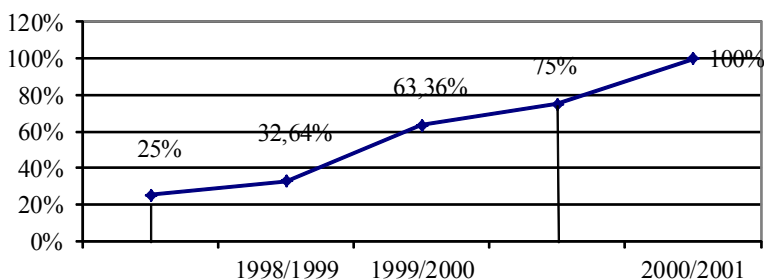


Рис. 2.4. Кумулята приема в первые классы общеобразовательных школ г. Харькова (отображение первого квартиля)

1998 года. Эта цифра означает, что 25% первоклассников были приняты в школы г. Харькова в 1998/99 учебном году. Второй и третий квартили находятся аналогично.

Именно легкость получения графического решения дает кумуляте некоторое преимущество перед таблицей кумулятивных частот, при пользовании которой избежать утомительных интерполяций невозможно.

Графики качественных данных. Графическое изображение качественных данных отличается определенным образом от графиков количественных данных. Для графического изображения качественных данных используется длина отрезка, площадь фигуры или интенсивность оттенка цвета. Здесь представляется только три простейших обычно встречаемых типа: 1) *диаграмма полос*; 2) *круговая (гартовская) диаграмма*; 3) *статистическая карта*.

Диаграмма полос. Диаграмма полос представляет собой последовательность равноотстоящих друг от друга полос, длины полос пропорциональны соответствующим им частотам (см. *Рис. 2.5*). Построим диаграмму полос для таблицы 2.5 (см. Стр. 66 данного учебника). Так как табличные признаки имеют только одно измерение, а именно – измерение частоты, то соответствующая диаграмма полос требует только одну шкалу – шкалу частот, которая обычно откладывается вдоль горизонтальной базовой линии. Так как в данном случае отсутствует

непрерывная количественная шкала (как это имеет место в гистограмме), которая определяет размер основания полосы, то полосы могут иметь любую удобную ширину и располагаться в любом возможном порядке. Они вычерчиваются одинаковой ширины, просто для наибольшей наглядности диаграммы; пустое пространство между ними – шириной в половину полосы – служит для подчеркивания дискретного характера качественных данных.

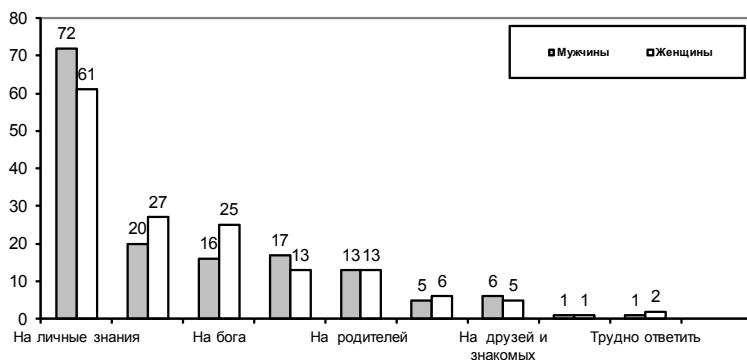


Рис.2.5. Диаграмма полос распределения ответов респондентов в зависимости от их пола на вопрос анкеты: «На кого/что Вы в наибольшей мере полагаетесь в своей жизни?» (% к опрошенным)

Круговая диаграмма – гартовская диаграмма. Как говорит само ее название, круговая диаграмма представляет собой круг, в котором вычерчиваются секторы, площадь которых пропорциональна наблюдаемым частотам (см. *Рис. 2.6*). При построении этой диаграммы частоты, выраженные в процентах, или частоты, последовательно вымеряются угломером вдоль окружности в 360° , или 100%. От этих отметок проводятся радиусы к центру круга, которые разделяют общую площадь на секторы, пропорциональные частотам. Наглядность круговой диаграммы, или, что то же самое, всех качественных диаграмм, часто может быть повышена путем раскрашивания. Построим круговую диаграмму для предыдущего примера (см. *Рис 2.5.*), отразив на ней распределение частот выбора мужчин (см. *Рис. 2.6*).



Рис. 2.6. Круговая (гартовская) диаграмма распределения ответов мужчин на вопрос анкеты: «На кого/что Вы больше всего полагаетесь в жизни?» (% ко всем опрошенным)

Статистическая карта. Географические изменения в таких вопросах, как состав населения, количество бракосочетаний и разводов, количество преступлений и экономических показателей, – обычный предмет социологического анализа. Систематическое выявление таких измерений может служить целям, как социальной политики, так и научного объяснения. Например, открытие относительно высокой смертности в определенных географических областях или концентрация правонарушений в конкретном городском районе составляет первый шаг в раскрытии причин этих явлений.

Общим принципом составления статистических карт является обозначение различных частот соответствующей плотностью штриховки, которая создается разнообразными поперечными штриховками или различной концентрацией точек. Существуют специальные справочники, в которых детально описываются стандартные способы графического представления, на которых не считаем необходимым останавливаться в данном издании.

Основной проблемой в географическом картографировании является ясное определение пространственной единицы измерения. Такие единицы измерения, как города или государства, в этом отношении обычно не создают трудностей, так как границы этих единиц зафиксированы по закону. Но такие единицы, как городские кварталы, национальные районы и естественные экономические области, ставят множество проблем, которые необходимо разрешить до составления карты. Несмотря на известные сложности, статистическая карта – важный инструмент социологического исследования. В частности, она является существенным элементом в экономической школе социологии, которая концентрирует внимание на пространственный аспект социального поведения. Статистическая карта состоит из тех контуров, которые считаются существенными для представления и понимания наносимых (на карту) социальных данных.

Временные диаграммы. Временные ряды представляются обыкновенно не в виде таблиц, а в виде графиков по той простой причине, что протяженные временные ряды трудно читать в виде таблиц, а исключительно короткие ряды, во всяком случае, не выразительны. Всякие отклонения от направления, скорости (флуктуации) можно намного легче обнаружить из графика, чем из ряда чисел.

Временные графики существуют *в двух преобладающих формах*: 1) *арифметические диаграммы* и 2) *полулогарифмические диаграммы*. В данном издании подробно остановимся только на арифметической диаграмме.

Арифметические временные диаграммы аналогичны простому графику частот, за исключением того, что строится график количественных наблюдений на интервалах времени, вместо того чтобы относить частоты событий к интервалам переменной. Интервалы времени наносятся вдоль горизонтальной оси, а флуктуации переменных величин наносятся на вертикальной оси.

Вместо ломаной линии временная диаграмма может состоять из ряда отдельных столбцов или полос, чьи высоты пропорциональны относительным величинам временных рядов. Эту временную диаграмму, которая является просто рядом календарных

наблюдений, не нужно путать с бинарным распределением таких величин, как возраст бракосочетания и размер семьи, в которых базовая линия представляет параметр времени. Когда время является явной переменной в таких бинарных распределениях, для него существует истинная нулевая точка, как при измерении возраста или в любом воспроизводимом измерении. Во временных рядах, однако, отметки на базовой линии весьма произвольно фиксируют отдельные моменты времени. Тем не менее такие моменты времени часто рассматриваются как интервал, чтобы упорядочить наблюдаемые данные.

Многозначный график. При некоторых обстоятельствах, возможно начертить на одной и той же арифметической шкале два или больше временных ряда (чтобы более ясно выявить соотношения между ними). Такой график вполне понятен, потому что две переменные имеют приблизительно одну и ту же область и общее расположение на шкале. Однако эта процедура должна производиться с некоторой предосторожностью, потому что она может ввести в заблуждение читателя в тех случаях, когда переменные располагаются неодинаково. Ложное впечатление возникает в результате того, что два ряда данных расположены на неравных расстояниях от начала отсчета. Чтобы облегчить сравнение степеней изменения строится полулогарифмический график, который определенным образом преобразует шкалу.

2.4. Характеристики положения **(среднее арифметическое, мода, медиана, квантили)**

Основные числовые характеристики одномерного распределения: максимум; минимум; средние величины.
Характеристики положения – величины, определяющие положение центра эмпирического распределения, устанавливающие положение всех единиц распределения вдоль шкалы. Определение этих характеристик может быть весьма результативным и в эмпирической социологии.

По окончании полевого этапа социологического исследова-

ния, как правило, социолог сталкивается с огромными массивами данных. Во всей своей полноте эти данные могут быть излишне подробными (дробными), слишком «неудобными» для обращения с ними, громоздкими для анализа и сравнения с другими аналогичными данными. Как мы уже говорили в начале данного раздела, для того чтобы преодолеть громоздкость и излишнюю дробность полученных данных, исследователь обращается к их агрегированию (укрупнению). В связи с этим он апеллирует к математической статистике, которая, в соответствии со своим назначением, должна: а) давать информацию о взаимном расположении фактов и событий; 2) исключать те величины, которые в данный момент не относятся к делу, 3) наглядно представлять общую картину.

Пределом эффективной агрегации, очевидно, было бы сведение множества событий к одной единственной величине, которая некоторым образом представляла бы собой их полную совокупность. Ясно, что никакая отдельная величина не может быть достаточно многосторонней, чтобы отразить каждую характеристику распределения; она может выразить только одно свойство из множества. Эта характерная величина не является точной копией общего. Она лишь приближенно описывает распределение значений.

Любая величина в распределении может описывать всю совокупность, если известно ее относительное положение в распределении. Следовательно, необходимо проанализировать все события в совокупности для того, чтобы оценить представительность любого из них. Но практически, не все величины одинаково полезны в этом отношении. Наиболее удобными и полезными для выделения их из таблиц являются: 1) максимум, 2) минимум и 3) центральные или типичные величины, известные как средние.

Максимум. В некоторых случаях максимальная величина переменной в распределении является единственной представительной величиной. При флуктуациях размера и веса движущегося транспорта такой величиной является наибольший груз, который можно безопасно перевозить по шоссе или мосту –

знания «среднего» совершенно недостаточно, так как мост, построенный для средней нагрузки, разрушился бы при максимальной.

Подобно этому вместимость школ, госпиталей и других институтов планируется для ожидаемого максимума, а не для предполагаемого среднего. Знание средней величины в данном случае бесполезно. Поэтому максимум, как мера положения, представляет собой граничную величину, ниже которой расположены не относящиеся к делу величины.

Минимум. Многие проблемы социологии политики сводятся к выбору минимальной величины распределения в качестве действующей нормы. Очевидно, что минимум – это такая величина, выше которой располагаются все другие величины в данном расположении. Так, для законодательства норм благосостояния населения из распределения доходов «нормальных» семей выводится минимальный доход (прожиточный минимум). Из гипотетического распределения зрелого населения по возрастам устанавливается минимальный возраст для женитьбы, службы в армии, права голосования и права избрания, а также и других социальных обязанностей. Будущий студент вуза может интересоваться только тем, каковы минимальные расходы, необходимые для обучения в нем.

Ни максимум, ни минимум не требуют сколько-нибудь серьезных вычислений как при определении их местоположения, так и в интерпретации. Их смысл весьма прост и, следовательно, не нуждается в более широком обсуждении.

Среднее, особенности выбора среднего. Наиболее общей и распространенной мерой расположения (и, в общем, наиболее полезной) является среднее. Обычно это центральная величина, вокруг которой группируется распределение. Явная тенденция многих статистических совокупностей концентрироваться вокруг центра часто называется «центральной тенденцией», а значение величины в этом центре – «мера центральной тенденции» – обычно называется средней.

Однако этот функциональный центр не обязательно идентичен с серединой области данных наблюдений. Область

наибольшей концентрации может находиться как вблизи средней точки области, так и на значительном расстоянии от нее. Распределение оценок в тестах на испытание умственных способностей имеет колоколообразную форму с максимумом в средней точке области. С другой стороны, например, распределение доходов часто представляется V-образными – кривыми, когда большинство единиц распределения сконцентрировано в левой и правой оси шкалы, а не посередине или J-образными – кривыми, центр которых существенно смещен в сторону одной из осей шкалы (см. *Рис 2.10*, с. 81).

Подобно большинству других статистических мер, понятие о среднем имеет свои корни в обычном здравом смысле. Каждый политический деятель, основываясь на результатах исследования «своего» электората, совершенно привычно говорит о «среднем избирателе», «средней семье», «среднем студенте» и т. п. Широкое разнообразие черт, которые сводимы к среднему, видно из популярного описания «среднего» мужчины Украины, «ростом в сто семьдесят девять сантиметров, весом 82 килограмма, который предпочитает брюнеток, футбол, сало, борщ и жаркое и считает, что способность вести домашнее хозяйство является наиболее важным достоинством жены». Не существует, разумеется, ни одного мужчины, который является средним во всех отношениях; человек среднего роста не обязательно будет человеком среднего ума или средней красоты. Поэтому популярное утверждение, что «среднего человека» не существует, полностью оправдано. Такие фиктивные понятия, как «средняя школа» и «средний украинец», не соответствуют более строгому статистическому понятию, которое применимо только к рядам измерений одной переменной. Тем не менее какой бы смысл не вкладывался в это понятие, политик бессознательно подразумевает то, что все статистики точно признают: среднее – есть разновидность нормы, вокруг которой колеблется переменная. Разница между политическим деятелем и социологом заключается в том, что последний требует большей точности, чем это позволяет неофициальное народное словоупотребление. Поэтому социолог-профессионал разрабатывает терминологию и математические

процедуры для измерения среднего и тем самым ограничивается переменными, у которых центральная тенденция допускает некоторые виды квантификации. Различным типам центральных тенденций соответствуют различные средние, каждое из которых отвечает требованиям данной проблемы. Из этих многочисленных типов средних в данном издании подробно обсуждаются только три: среднее арифметическое, мода и медиана.

Среднее арифметическое ($M[X]$). Как и в обычном словопотреблении, на статистическом языке «среднее» означает «типичное», «обычное», «ожидаемое». Среднее арифметическое – такое значение признака, сумма отклонений от которого всех значений признака равна нулю (с учетом знака отклонения). Если речь идет о выборочном исследовании, в основании которого лежит случайная выборка, среднее арифметическое часто называется математическим ожиданием, так как в данном случае имеется в виду среднее значение случайной величины. Следует подчеркнуть, что условное обозначение среднего в данном случае может отличаться. Как правило, если имеется в виду среднее по выборке, то вместо $M[X]$ используются $\bar{x}$, что необходимо запомнить, во избежание путаницы в формулах. Считается, что любое физическое тело, находясь в неопределенном состоянии, будет стремиться принять состояние равновесия (опоры на свой центр тяжести), точно так же и среднее значение любой случайной величины при достаточно большом количестве испытаний будет стремиться к своему математическому ожиданию. Этот факт доказывается в теории вероятности. Математическим ожиданием случайной величины называется сумма произведений всех возможных значений случайной величины на вероятности этих значений.

В целом вычисление среднего арифметического необходимо для осуществления более сложных математических операций, связанных с анализом социологических данных, в чем можно будет убедиться на последующих занятиях. Именно среднее арифметическое, как математическое ожидание, является одной из основных характеристик выборки, так как с его помощью можно прогнозировать значения некоторого случайного признака

при достаточно долгом периоде испытаний. Например, человек, удрученный затянувшейся «полосой неудач», обычно надеется, что события должны прийти в норму. Он полагает, что существует нечто вроде закона природы, согласно которому полоса неудач должна сбалансироваться удачами.

Процедуры вычисления среднего арифметического для несгруппированных и сгруппированных данных несколько отличаются.

Вычисление среднего арифметического для несгруппированных данных. Для несгруппированных данных среднее вычисляется по простой формуле, путем суммирования отдельных величин и последующего деления на общее число событий, что имеет следующий вид:

$$M[X] = \frac{\sum x_i}{N} \text{ (формула простого среднего),}$$

где • $M[X]$ – среднее арифметическое;

- $\sum x_i$ – сумма переменных;
- x_i – значение переменной (ее величина);
- N – число событий.

Вычисление среднего арифметического для сгруппированных данных. В первом подразделе этого раздела мы говорили о том, что сгруппированные данные отличаются от несгруппированных тем, что каждой группе подобных величин приписывается частота или «вес». Поэтому для вычисления среднего сгруппированных событий каждое значение переменной x_i умножается на свою частоту (n_i), затем эти произведения суммируются, и вся сумма делится на сумму всех частот, что имеет следующий вид:

$$M[X] = \frac{\sum x_i \cdot n_i}{N} \text{ (формула среднего взвешенного),}$$

где • x_i – значение переменной (ее величина);

- n_i – частота встречаемого значения переменной;
- N – сумма всех частот, объем исследуемой совокупности.

В этой связи важно понимать, что если применить две разные формулы (простого и взвешенного среднего) к одному и тому же распределению, значения средних, полученные по двум разным формулам, отличаться не будут.

Вычисление среднего арифметического для интервальных рядов. Процедура вычисления среднего арифметического для интервальных рядов данных, несколько отличается от процедур, описанных выше. Необходимо понять, что для непрерывных интервальных рядов частота интервала совпадает с его средней точкой. Следовательно, перед тем как вычислять среднее арифметическое для интервального ряда, необходимо найти средние точки каждого интервала. Для этого сначала находятся границы интервалов, а затем вычисляются их средние точки по формуле:

$$x_{ci} = \frac{(x_i + x_{i+1})}{2},$$

где • x_i – нижняя граница интервала;

• x_{i+1} – верхняя граница интервала.

После чего каждая средняя точка взвешивается (то есть перемножается на частоту интервала) и вычисляется среднее арифметическое для всего интервального ряда по формуле:

$$M[X] = \frac{\sum x_{ci} \cdot n_i}{N},$$

где • x_{ci} – средняя точка интервала;

• n_i – частота интервала;

• N – сумма всех частот интервального ряда.

Вычисление среднего ряда средних (комбинированное среднее). Две или более средние величины сами часто усредняются, т. е. можно получать среднее ряда средних. Средние подгрупп до того, как они будут скомбинированы, должны быть взвешены в соответствии со своими N . Недооценка необходимости взвешивать средние может привести к абсурдным результатам. Можно привести пример из игры в баскетбол. Игрок за 75 ударов набрал 25 очков, что дало в среднем 0,33. В этот же день он за пять ударов набрал пять очков, что дало в среднем – 1,00. Каково будет среднее (комбинированное)?

Наивным было бы предположение о том, что если мы сложим две средние величины и разделим затем на 2 (по принципу нахождения среднего арифметического простого), то получим искомый результат. В таком случае, по отношению к рассматриваемому примеру, комбинированное среднее будет равным 0,667 – явно неправильный результат. В таблице 2.7 проиллюстрирована процедура правильного вычисления комбинированного среднего для приведенного выше примера.

Таблица 2.7

Расчетная таблица по нахождению комбинированного среднего

<i>Количество ударов (n_i)</i>	<i>Общее количество очков (Σx_i)</i>	<i>Среднее количество очков за каждый удар ($\bar{x}_i$)</i>	$\bar{x}_i \times n_i$	<i>Конечный результат – $M[X]$</i>
75	25	25/75=0,33	75×0,33=24,75	
5	5	5/5=1,00	5×1,00=5,00	
ИТОГО: $N=80$			24,75+5,00=29,75	$M[X] =$ $=29,75:80=$ $=0,37$

Получается, что в среднем за каждый удар в этот день баскетболист получил приблизительно по 0,37 баллов (число, округленное до сотых).

Приведем более близкий социологической практике пример нахождения среднего арифметического для ряда средних. Предположим, нам известны средние показатели численности детей, воспитывающихся в дошкольных учреждениях каждого из районов города. Но перед нами стоит задача выйти на один средний показатель по всему городу. Порядок вычислений в направлении разрешения этой задачи проиллюстрирован в таблице 2.8.

Для определения среднего количества детей ($M[X]$) в дошкольных учреждениях г. Харькова в 1999 г. необходимо первым делом узнать общее количество детей ($N = \sum \bar{x} \times n = 33300$),

Таблица 2.8

**Сеть дошкольных учреждений всех ведомств
г. Харькова в 1999 г.¹³**

№ п/п	Название района	Количество д/у (n_i)	В них детей в среднем ($\bar{x}_i$)	$\bar{x}_i \times n_i$
1	Дзержинский	33	161,06	33×161,06=5315
2	Октябрьский	13	140,62	1828
3	Киевский	33	141,88	4682
4	Коминтерновский	17	182,82	3108
5	Ленинский	23	93,09	2141
6	Московский	35	208,46	7296
7	Орджоникидзевский	24	151,88	3645
8	Червонозаводский	18	126,00	2268
9	Фрунзенский	16	188,56	3017
10	Итого (N)	212		33300

а затем полученный результат разделить на количество дошкольных учреждений – 212. Продолжая этот процесс, найдем, что $M[X] = 157,07$, то есть в среднем на одно дошкольное учреждение г. Харькова в 1999 г. приходится 157,07 ребенка. Дробный характер дискретной величины (в реальной жизни не бывает 1,5 человека и т. п.), обычно, никого не смущает, если речь идет о среднем показателе.

Подводя черту под всем сказанным в отношении среднего арифметического, выделим его *основные свойства*:

- 1) сумма отклонений различных значений признака от среднего арифметического равна нулю;
- 2) если от каждой варианты вычесть или к каждой варианте прибавить какое-либо произвольное постоянное число, то среднее увеличится или уменьшится на то же самое число;

¹³ Стат. збірник: показники роботи закл. освіти та наук. установ обл. за 1999 рік [за заг. редакцією О. Л. Сидоренка, А. С. Доценка, П. С. Дементьєва]. – Х., 2000. – 82 с.

3) если каждую варианту умножить (разделить) на какое-либо произвольное постоянное число, то среднее увеличится (уменьшится) во столько же раз;

4) если веса, или частоты, разделить или умножить на какое-либо произвольное постоянное число, то величина среднего не изменится.

Необходимо уяснить и запомнить, что *среднее арифметическое вычисляется и имеет смысл только для метрических и порядковых шкал*. Оно представляет величину каждого события в распределении. В связи с этим оно подвержено влиянию как очень больших, так и крайне малых величин, что особенно заметно в несимметричных распределениях. Для таких распределений более информативными могут быть иные меры усреднения, такие как мода и медиана.

Мода или вероятностное среднее. *Мода представляет собой наиболее часто повторяющуюся величину в упорядоченном распределении; она характеризует то место распределения, где концентрация событий максимальна.* Этимологически она связана с представлением о преобладающей манере одеваться или с этикетом, к которому, как предполагается, приспосабливается большинство той или иной социальной группы. Следовательно, моду (M_0) можно также определить как наиболее вероятную величину, и поэтому ее называют вероятностным средним.

Исследование разговорных выражений наводит на мысль, что под модой в действительности часто подразумевают понятие «среднее». Отчасти это объясняется тем, что признаки, так же как и переменные, могут иметь преобладающие частоты в серии наблюдений. Когда политический деятель говорит о «среднем избирателе», интересы которого он собирается отстаивать, он исходит из того, что большинство избирателей руководствуются своими собственными интересами; или когда официантка замечает, что «средний» клиент не пьет черный кофе, она, вероятно, имеет в виду, что большинство посетителей данного ресторана не пьют черного кофе. Посетители ресторана и избиратели в свою очередь говорят о «средней официантке» и «среднем политическом деятеле».

На языке статистики *мода* – это величина, с которой наиболее вероятно можно встретиться в серии зарегистрированных наблюдений. Таким образом, мода является наиболее вероятной величиной, хотя знание одной лишь моды не позволяет определить степень этой вероятности. Понятно, что если бы частоты всех величин были одинаковыми, то не было бы никакого смысла вводить это понятие.

Вычисление моды. В преобладающем большинстве случаев для определения моды достаточно просто подсчитать частоту появления каждой величины, так как мода – наиболее часто наблюдающаяся величина. В случае непрерывных данных подразумевается, что эмпирические измерения настолько подробны и кратковременны, что никаких два измерения не дают тождественных результатов. Очевидно, что мода не может выявиться без группировки, так как в противном случае не было бы никаких частот. Само собой разумеется, в случае наличия интервалов они должны быть одинаковыми по ширине; если это правило не соблюдать, можно, выбирая достаточно большие интервалы, получить моду практически любой желаемой величины – явно бессмысленный результат.

Следовательно, необходимо пройти две различные стадии в определении моды: 1) определение модального интервала и места в нем преобладающей или «модальной» частоты и 2) нахождение величины, соответствующей этой частоте. В качестве примера рассмотрим таблицу 2.9.

В таблице 2.9 наибольшая частота равна 73, следовательно, интервал 1994/1995 гг. является модальным, а модальная величина или приближенная мода равна 1994,5, что является средней точкой модального интервала.

Однако существование «произвола» в выборе ширины интервала вносит некоторую неопределенность в процедуру вычисления моды. Существует два возможных пути избегания этой неопределенности: 1) отказаться от данного типа средней величины и применить другую меру среднего; 2) применить вместо приближенного более точный метод вычисления, который уменьшил бы неопределенность. Следует отметить, что даже при

Таблица 2.9

Количество высших учебных заведений Украины III–IV уровня аккредитации, появившихся в определенный учебный год
(известно, что в 1992 году их было 1580)

Годы ($x_i; x_{i+1}$)	Количество учеб. завед. III–IV уровня аккредитации (n_i)	Накопленные частоты (n_i^H)
1993–1994	1	1
1994–1995 (Модальный интервал)	73	74
1995–1996 (Медианный интервал)	23	97
1996–1997	19	116
1997–1998	6	122
1998–1999	18	140
1999–2000	15	155
<i>Итого</i>	<i>155</i>	

условии неопределенности в вычислении моды довольно часто в социологической практике бывают такие случаи, когда нельзя отказаться от моды в пользу других мер усреднения. Следовательно, нужно обратить внимание на усовершенствование техники вычисления.

В приближенном методе, рассмотренном выше, находится средняя точка интервала с наибольшей частотой; при этом игнорируются смежные интервалы и их частоты. Однако эти смежные интервалы повлияли бы на величину моды, если бы их границы были иначе расположены. Следовательно, необходимо разработать более «чувствительный» метод, который позволит учесть влияние смежных интервалов.

Метод разностей в вычислении моды. Более совершенный, чем вышеописанный, метод вычисления моды сводится к следующему: 1) вычисляются разности между модальной частотой и частотами смежных интервалов; 2) вычисляется отношение одной из этих разностей (обычно с частотой нижнего интервала) к сумме двух других разностей; 3) это отношение умножается

на ширину модального интервала, а затем полученный результат прибавляется к истинной нижней границе модального интервала. В итоге получаем уточненное значение моды. Вычислительная формула в данном случае будет иметь следующий вид:

$$M_o = x_0 + h \frac{n_{mo} - n^-}{2n_{mo} - n^+ - n^-},$$

где • x_0 – нижняя граница модального интервала,

• h – ширина модального интервала,

• n_{mo} – частота модального интервала,

• n^- – частота интервала, предшествующего модальному интервалу,

• n^+ – частота последующего интервала.

Применив формулу к таблице 2.9, получим следующие результаты:

$$x_0 = 1994 \quad h = 1 \quad n_{mo} = 73 \quad n^- = 1 \quad n^+ = 23$$

$$M_o = 1994 + 1 \times \frac{73 - 1}{2 \cdot 73 - 1 - 23} = 1994,59.$$

Бимодальность. Некоторые распределения обнаруживают два максимума и поэтому называются бимодальными, в отличие от унимодальных распределений. Бимодальность распределения может быть следствием наложения двух или более популяций с различными частотными максимумами. Так, например, в полигоне распределения роста взрослого населения, благодаря объединению групп мужчин и женщин, которые характеризуются двумя различными распределениями роста, может возникнуть бимодальность. Столкнувшись с бимодальностью, исследователь должен попытаться либо разделить распределения, которые ее вызывают, либо (в случае неудачи) принять бимодальность как характеристику, присущую данному распределению.

Медиана. В любой упорядоченной совокупности каждое событие занимает определенное место – первое, второе, десятое или семьдесят пятое или ранг. Очевидно, что любая конкретная численная величина ранга приобретает смысл и значение

в зависимости от общего числа рангов. Ранг, равный 10, в ряду из 100, является относительно более высоким, чем ранг 10 в группе из 20.

Точка, которая рассекает упорядоченную (ранжированную) совокупность на две равные части так, что одна половина событий точно находится ниже, а другая половина выше этой точки, называется медианой. Например, в 1950 г. медианный возраст всего населения Соединенных Штатов был равен 30,4 года. Это означает, что одна половина населения была старше, а другая половина населения моложе этого возраста. Так как медиана ясно обозначает положение величины в последовательности, она часто называется «средним положением».

Вычисление медианы. По аналогии со средним арифметическим, существуют разные формулы вычисления медианы для несгруппированных и сгруппированных данных.

Если данные *не сгруппированы*, как правило, придерживаются одной из следующих формул:

$$\frac{N}{2} \text{ или } \frac{N+1}{2},$$

где N – общее число рангов. При этом, первая формула чаще используется при нечетном количестве рангов, а вторая – при четном количестве.

В случае *если данные сгруппированы*, – вычисления осуществляются по следующей формуле:

$$Me = x_0 + h \frac{\frac{1}{2}N - n_H}{n_{me}},$$

где • x_0 – нижняя граница медианного интервала;

• h – ширина медианного интервала;

• N – объем выборки (соответственно, $\frac{1}{2}N$ – «половинное»

событие);

• n_H – частота, накопленная до медианного интервала;

• n_{me} – частота медианного интервала.

Если вкратце описать алгоритм вычисления медианы для данных, сгруппированных в интервалы, то следует начать с (1) ранжирования всех событий или нахождения накопленных частот для каждого интервала, затем (2) найти, так называемое, «половинное событие» (ведь медиана – точка, делящая пополам), (3) по ряду накопленных частот найти интервал, в который попадает это событие, и (4) осуществив элементарные математические вычисления, найти саму медиану.

Применим данный алгоритм вычисления медианы к таблице 2.9, представленной выше: (1) ранжируем все события путем нахождения накопленных частот для каждого интервала, (2) делим сумму всех частот пополам: $N/2=155/2=77,5$ («половинное» событие), (3) находим медианный интервал, то есть местоположение 77,5-го события в ряду накопленных частот (*очевидно, что «половинное» событие не попадает в первые два интервала, накопленные частоты которых меньше, чем 77,5; частота третьего интервала, равная 97, существенно превосходит отмеченную границу*), получаем, что 77,5-е событие находится где-то внутри интервала 1995/1996 гг. с частотой в 23, которая, по предположению, однородно распределена по всему интервалу.

Теперь можно определить все неизвестные компоненты формулы вычисления медианы: $x_0 = 1995$, $h = 1$, $\frac{1}{2}N = 77,5$, $n_H = 74$, $n_{me} = 23$.

Полученные значения подставим в соответствующую формулу:

$$Me = 1995 + 1 \times \frac{77,5 - 24}{23} = 1995,15.$$

Интерпретируя эту цифру, скажем, что ровно половина всех учебных заведений III–IV уровня аккредитации появилась в Украине до момента 1995,15 года, и ровно половина – позже.

Медиана дискретных данных. Некоторые авторы ограничивают применение медианы только непрерывными данными по той причине, что дискретные данные по определению не могут быть дробными, как это требуется для медианы. Но такое

ограничение не представляется столь необходимым. В распределении размеров семей нет такого размера семьи, чтобы точно 50% семей были больше, а 50% семей – меньше. Можно ли в такой ситуации отказаться от медианы или следует прагматически трактовать данные как непрерывные и принять дробную величину в качестве медианы? Как и ранее, примем последнюю альтернативу. Ведь сущность медианы крайне проста: 50% событий имеют меньшую величину и 50% – большую. Она обладает также еще одним показательным критерием: суммарное расстояние между медианой и каждой из величин распределения всегда меньше, чем подобная величина, вычисленная для любой другой точки. Именно поэтому медиана «ближе» к событиям данного распределения, чем любая другая мера среднего. В этом и состоит смысл того, что медиана занимает центральное положение в распределении.

Другие меры усреднения. Вместо того чтобы вычислять медиану, можно было бы для большей точности локализовать данные в меньшем интервале – верхней четверти, десятой или даже сотой. Для таких дополнительных уточнений требуются меньшие подразделения. Вычисление квартилей, децилей и центилей, которые разделяют множество событий на четвертые, десятые и сотые части, выполняются совершенно аналогично вычислению медианы, за исключением того, что в формулу для медианы, данную выше, подставляют вместо другую величину, которая соответствует частоте или процентам событий, лежащих ниже рассматриваемого места. Например, для нахождения точки, ниже которой приходится наименьшая четверть случаев, заменяют $N/2$ на $N/4$.

Таким образом, первый квартиль в распределении будет равен:

$$Q_1 = x_o + \delta \frac{N/4 - n_{H1}}{n_{Q1}}.$$

Если потребуется выделить точку, ниже которой находятся

75% событий (или Q_3), то вычисления производятся по следующей формуле:

$$Q_3 = x_o + \delta \frac{3N/4 - n_{H3}}{n_{Q3}}.$$

90-й центиль (или C_{90}) может быть найден по формуле:

$$C_{90} = x_o + \delta \frac{90N/100 - n_{H90}}{n_{C90}}.$$

Квантили в качестве нормирующего критерия. Медиана, квартили, квантили, децили и центили, которые, согласно своему определению, указывают на долю событий, расположенных ниже или выше данной величины, носят обобщенное название *квантили*. Эти меры усреднения можно применять для фиксации относительного положения любой данной величины в своем ряду. Можно локализовать с помощью 90-го центиля вес в 80 килограмм ($C_{90}=180$). Это будет означать, что 10% населения обладают весом выше, а 90% – меньше данного веса. Точно так же на абстрактной шкале центилей можно локализовать и такие события, как возраст в 62 года, рост в 180 см, уровень интеллекта в 120 баллов и т. д.

Очевидно, что квантили можно рассматривать как стандартизованные меры расположения независимо от метрической системы или типа данных. Таким образом, человек, который находится на 90-м центиле в умственном развитии, может находиться вблизи 90-го центиля и по доходам, что указывает на сходство между этими двумя социальными явлениями. Такой подход позволяет с успехом сопоставлять и сравнивать «несравнимые» величины. Данные характеристики положения не зависят также от вида распределения, то есть от того, является ли оно нормальным, скошенным или прямоугольным. Это обстоятельство еще более повышает ценность квантилей, позволяя с помощью вышеописанных процедур ранжировать любую случайную величину в некотором упорядоченном множестве.

Подводя итог всему сказанному в данном параграфе, представим краткое изложение характеристик средних:

Среднее арифметическое

1. Это такая величина в данной совокупности, которая наблюдалась бы в том случае, если бы все величины были равны.
2. Суммы отклонений от среднего в любую сторону равны; следовательно, алгебраическая сумма отклонений равна нулю.
3. Среднее арифметическое репрезентирует значения каждой величины распределения.
4. Множества имеют одно и только одно среднее.
5. Над средним можно производить алгебраические действия, средние подгрупп можно комбинировать при надлежащем взвешивании.
6. Среднее можно вычислить даже в том случае, когда значения отдельных величин неизвестны при условии, что известна сумма всех величин (N).
7. Для вычисления среднего нет никакой необходимости группировать и упорядочивать величины.
8. Среднее можно вычислить только для закрытых интервалов.
9. Оно фиксировано в том смысле, что процедуры группирования не оказывают сколько-нибудь серьезного влияния на нее.
10. Среднее применимо только к количественным данным.

Мода

1. Наиболее часто наблюдаемая величина в распределении; точка наибольшей плотности.
2. Значение моды определяется преобладающей частотой, а не значениями переменной в распределении.
3. Это наиболее вероятная величина и, следовательно, наиболее типичная.
4. Данное распределение может иметь две или более моды. С другой стороны, в прямоугольном распределении не существует никакой моды.
5. Мода не выражает степени модальности.

6. Над модой нельзя производить алгебраические манипуляции; моды подгрупп нельзя комбинировать.

7. Она неопределенная в том смысле, что зависит от процедуры группировки.

8. Мода определена как для открытых, так и для закрытых распределений.

9. Мода – единственный тип среднего, который может представлять качественные переменные.

Медиана

1. Это величина, расположенная точно в средней точке множества (а не в области изменения переменной); половина событий имеет значения больше ее, а половина – меньше.

2. Значение медианы определяется ее расположением во множестве данных и не зависит от значения отдельных величин.

3. Сумма расстояний от медианы до других величин множества меньше, чем подобная сумма, вычисленная для любой другой точки распределения.

4. Каждое множество имеет одну и только одну медиану.

5. Над медианой нельзя производить алгебраические действия: медианы подгрупп не могут взвешиваться и комбинироваться.

6. Она фиксирована в том смысле, что процедура группирования не оказывает на нее заметного влияния.

7. Для вычисления медианы все величины должны быть упорядочены и сгруппированы.

8. Медиана вычисляется как для открытых, так и для закрытых интервалов.

9. Качественные данные не позволяют рассчитать медиану.

Подчеркнем, что отнюдь не во всех случаях является целесообразным определение всех характеристик положения. Существует ряд критериев, помогающих решить вопрос о применимости того или иного типа среднего. Об этих критериях и пойдет речь в следующем подразделе.

2.5. Критерии выбора вида усреднения

Сопоставимость средних. В связи с тем что нам часто приходится выбирать между различными видами усреднения, необходимо изложить основные принципы выбора центрального значения. Уже было выяснено, что не существует универсальных типов среднего, которые можно применять во всех случаях. Необходимо всегда помнить, что среднее является единственным представителем распределения, величиной, которая весьма удобна вследствие компактности; однако она в то же время неудобна из-за краткосрочности. В лучшем случае, среднее вскрывает и вычеркивает столько же информации, сколько извлекает из распределения.

Известно три критерия, которые помогают решить вопрос о применимости того или иного типа среднего: (1) цель усреднения, (2) вид распределения данных, (3) ограничения, связанные с «техническими» причинами и типом измерительных шкал.

Цели усреднения. Любое эмпирическое социологическое исследование по существу является попыткой дать ответ на вопрос о природе явления, и статистическая процедура выступает лишь в роли инструмента. Исследователя при вычислении того или иного типа среднего интересуют следующие вопросы: каковы размеры семьи, какова продолжительность жизни, каков возраст населения.

Если размер семьи необходим для целей планирования жилищного строительства, то приближенная мода была бы более подходящей величиной, чем арифметическое среднее или медиана, даже если точная степень модальности неизвестна. Дома строятся не для абстрактных арифметических средних семей, а для реально существующих. Если же необходимо изучить плодовитость, то среднее арифметическое было бы более полезно, так как оно представляет как большие, так и малые семьи.

Вид распределения. Распределения могут иметь самый различный вид, от идеально симметричного до крайне асимметричного. Симметрия означает, что величины распределены

идентично по обе стороны от среднего. Степень асимметричности влияет на типичность и представительность средних величин, и, следовательно, ее необходимо принимать во внимание при выборе типа среднего. Иногда утверждают, что если кривая симметрична, то вообще не возникает никаких проблем в выборе способа усреднения, так как среднее, медиана и мода становятся тождественными. Однако это справедливо только в арифметическом смысле, но не в теоретическом. Даже если численные величины средних тождественны каждому типу – среднему арифметическому, моде, медиане, то им при этом соответствуют совершенно различные «образы». Например, «средний» студент считает свою оценку не суммой всех оценок, деленной на N , а скорее просто типичной оценкой. Эти примеры приводят к выводу, что даже при симметричном распределении выбирают такой тип среднего, который, прежде всего, удовлетворяет целям исследования.

По мере того как распределение становится все более и более асимметричным, величины разных типов среднего начинают заметно отличаться, и проблема выбора способа усреднения приобретает серьезное значение. Прежде всего, среднее арифметическое для несимметричного распределения перестает быть типичным. В U -образном распределении среднего вообще практически может и не быть, и поэтому его величина может быть совершенно фиктивной. Многие социологические данные: заработные платы, размеры городов, размеры семей и т. п. – часто несколько асимметричны, что требует особого внимания при выборе типа среднего.

Ограничения, связанные с «техническими» причинами и типом измерительных шкал. Существуют определенные чисто технические особенности вычислений, которые могут вынудить использовать тот или иной тип среднего. Так, например, арифметическое среднее нельзя вычислить для распределений с открытыми интервалами. В этом случае пользуются медианой, если распределение не очень асимметрично, то есть когда среднее и медиана не на много отличаются друг от друга. С другой стороны, когда известны только сумма частот и величин,

можно вычислить только среднее арифметическое, хотя другие типы среднего были бы предпочтительнее.

Необходимо всегда принимать во внимание технические возможности вычислительных методов. Так как не существует никакого метода комбинирования или взвешивания медиан и мод, то они обычно являются завершающим этапом вычислений. Поэтому, когда предвидятся дополнительные вычисления, необходимо выбирать среднее арифметическое и его производные.

Кроме того, выбор той или иной меры усреднения может ограничиваться типом шкалы, с помощью которой измерялся признак. Об этом мы подробно рассказывали в одном из предыдущих параграфов, посвященных рассмотрению различных типов шкал, применяемых в социологическом исследовании. Вспомним, что наибольшие ограничения в этом плане накладывает номинальная шкала. Распределение, полученное по номинальной шкале, мы можем охарактеризовать только с помощью моды.

Минимум, максимум и промежуточные меры. Во многих случаях только средние величины рассматриваются в качестве мер расположения. Однако необходимо подчеркнуть, что данная точка зрения слишком ортодоксальна. Мера расположения не обязательно должна совпадать с мерой типичности. Средние величины, как правило, выражают типичность или представительность; но точно так же максимум, минимум или любая промежуточная величина могут служить мерой расположения, если только они соответствуют всему распределению.

Характеристики средних. Предпосылкой для надлежащего применения каждого типа среднего является знание соответствующих описательных характеристик. Последние часто анализируются с точки их «преимуществ и недостатков». Но такой подход является скорее оценочным, чем описательным, и поэтому не дает строгого изложения сущности каждой процедуры. Преимущества при решении одной проблемы могут оказаться недостатками при решении другой. Поэтому можно формально изложить описательные характеристики, присущие каждому типу среднего, независимо от ситуации, в которой они могут применяться.

Сравнение типов среднего. Подобно любому статистическому показателю средние, прежде всего, применяются для целей сравнения. Наиболее часто средние величины используются для сравнения положений отдельных групп и распределений на одной и той же шкале. Очевидно, однако, что различные типы средних величин принципиально не сравнимаемы. Например, среднее арифметическое одного распределения нельзя сравнивать с модой другого распределения.

Различие характеристик средних становится еще более очевидным на примере асимметричных распределений. Будучи подверженным влиянию каждой величины, среднее арифметическое асимметричного распределения будет смещаться в направлении экстремальных величин. На моду края распределения не оказывают никакого влияния, в то время как медиана смещается преимущественно по направлению к «хвосту» асимметричного распределения. Однако это смещение не очень велико, поскольку концентрация событий в «хвосте», как правило, невелика. Если одномодальное распределение скошено вправо, порядок расположения различных типов среднего на базовой линии будет таков: мода, медиана, среднее арифметическое и интервалы между ними будут изменяться в зависимости от степени искажения (см. Рис. 2.7). Если распределение скошено влево, порядок средних будет обратным.

Но даже однотипные средние можно сравнивать только при условии подобия распределений, в том случае, когда «прочие

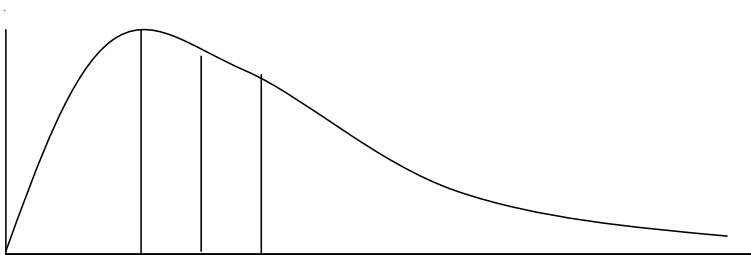


Рис. 2.7. Среднее, медиана и мода: правая скошенность

условия равны». Например, сравнение среднего роста мужчин и среднего роста женщин вполне обосновано. Степень различия среднего роста мужчин и женщин надлежащим образом измеряется разницей их средних арифметических. Однако когда распределения заметно отличаются друг от друга, такие сравнения фактически могут ввести в заблуждение; особенно в случае сравнения средних арифметических.

Например, два ряда величин могут иметь равные средние арифметические и совершенно различные типы распределения (как показано на *рис. 2.8*).

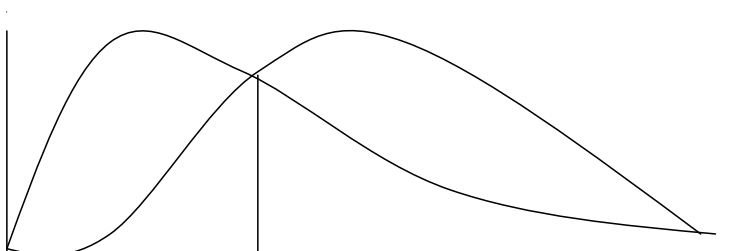


Рис. 2.8. Противоположная скошенность кривых.
Одинаковые средние

Эти распределения занимают одинаковую область. Однако их максимумы не совпадают, и поэтому было бы более целесообразно сравнивать данные распределения по их модам.

На рисунке 2.9 максимумы расположены приблизительно в одной и той же области; однако средние арифметические не

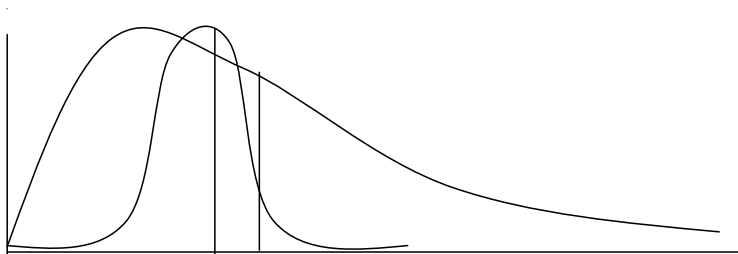


Рис. 2.9. Влияние несимметричности распределения на среднее

равны из-за несимметричности (скошенности) одного из распределений.

Делая общий вывод, следует подчеркнуть, что основные характеристики положения, о которых шла речь в этом параграфе, не представляют полного и всестороннего описания данных. Их следует рассматривать не как автономные величины, а лишь как определенные способы репрезентации конкретных данных. Использование средних ставит важную проблему реконструкции характера исходного распределения по нескольким извлеченным из него величинам. При решении этой проблемы важно знать другие характеристики – рассеяния, к рассмотрению которых мы и предлагаем обратиться в следующем подразделе учебника.

2.6. Характеристики рассеяния (дисперсия, отклонение, коэффициент вариации)

Меры вариации (протяженности) признака. Как уже говорилось, характеристики положения, хоть и являются чрезвычайно важными при изучении варьирующего признака, все же не дают полной информации о нем. Нетрудно представить себе два эмпирических распределения, у которых средние одинаковы, но при этом у одного из них значения признака рассеяны в узком диапазоне вокруг среднего, а у другого – в широком. Поэтому наряду с характеристиками положения нередко определяются и характеристики рассеяния исследуемой совокупности, показывающие, насколько близко/далеко все значения признака отдалены от средних показателей. Характеристики рассеяния выражаются в мерах вариации или мерах протяженности, наиболее употребляемыми из которых являются дисперсия, отклонения (среднее линейное и среднее квадратическое), коэффициенты вариации (для линейного и квадратического отклонений).

Вообще, понятие вариации лежит в основе всех статистических расчетов. Многие учебники не устанавливают различия между терминами «варьируемость» и «вариация». Однако заострим внимание на этих различиях. **Варьируемость** – способ-

ность изменяться. **Вариация** – проявление такой способности, которую можно описать и измерить.

Если бы все величины исследуемого множества были идентичными, то вычисление среднего значения или любой другой статистической величины стало бы излишним. Ведь основная цель усреднения – получение одной величины, которая бы репрезентировала (представляла) целую группу неодинаковых величин. Средние величины и были изобретены для исключения различий между величинами. Однако при определенных обстоятельствах для социолога представляет больший интерес характер отклонений, нежели результат усреднения. Например, оценивая общие способности и успеваемость студента, преподаватель должен принимать во внимание не только его средний балл, но и тенденцию этих баллов. Студент с баллами «5», «4» и «3», в общем, будет оценен иначе, чем студент с баллами «4», «4», «4», хотя средний балл у них одинаковый – «4». Тренеру по баскетболу, который отбирает для университетского первенства одного из двух игроков, имеющих равный средний балл, больше подходит стабильный игрок, который редко отклоняется от среднего, нежели неустойчивый спортсмен, который в среднем показывает низкий уровень игры и лишь иногда эффектно проявляет себя. Аналогично, в двух профессиональных группах, имеющих приблизительно одинаковый среднегодовой заработок, например, профессоров и бизнесменов, могут быть представлены самые различные заработки. Среди профессоров университетов оклады сравнительно мало отклоняются от среднего, тогда как в бизнесе доход является менее устойчивым.

В общем случае знание *картины вариации наблюдаемых значений* – того, что статистики называют *разбросы, рассеиванием или дисперсией*, не менее важно для социолога, чем знание средних величин. Для изучения этой картины в арсенале современной статистики имеется довольно много испытанных приемов и методов. Хотя они различаются в деталях, их можно разделить на три широкие категории в соответствии с процедурой их применения: 1) измерение области, содержащей всю или основную часть распределения; 2) измерение отклонений

переменной от центрального значения; и 3) измерение степени однородности качественных переменных.

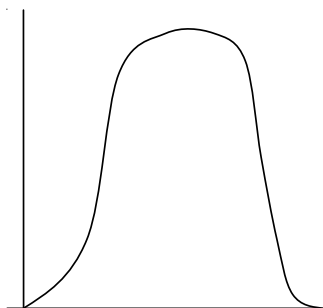
Некоторое общее представление о дисперсии количественных данных можно получить из графиков, представленных на рисунке 2.10. График частот сразу же позволяет прийти к выводу о том, является ли дисперсия симметричной или нет, увеличиваются ли частоты при приближении к среднему арифметическому (унимодальное распределение) или же они возрастают при смещении к концам интервала (*U*-образное распределение); распределяются ли величины равномерно по всей области вариации переменной или же они концентрируются в двух точках, как, например, в бимодальном распределении.

Однако такие визуальные впечатления являются индивидуальными и субъективными, а выводы о дисперсии, сделанные на их основе, вряд ли могут претендовать на научность. В таком случае научная обоснованность информации обеспечивается за счет объективных показателей, отражающих ту или иную меру вариации (протяженности, разброса, рассеяния), которые находят с помощью стандартных вычислительных процедур

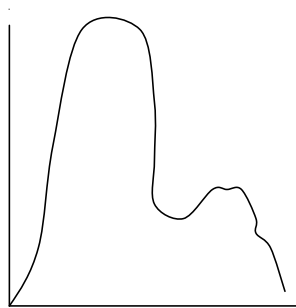
Измерение размаха вариации. Простейшая и самая грубая мера вариации – размах вариации (или «диапазон») *как размер области вариации переменной. По определению, **размах вариации** – это интервал, заключающий в себе все значения.* Следовательно, он находится точно так же, как и обычный интервал: вычисляется разность между истинными крайними значениями множества переменных, устанавливающими границы размаха вариации:

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

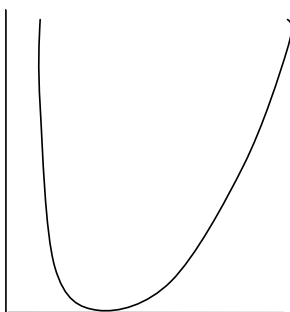
Например, для определения диапазона количества студентов высших учебных заведений III–IV уровня аккредитации нужно найти экстремальные истинные значения, а затем вычитаем одно из другого. Если наименьшее значение 159, а наибольшее – 259, диапазон будет равен разности между этими числами, то есть – 100. Таково расстояние, необходимое для расположения всех наблюдаемых частот. Для другого множества частот, несомненно,



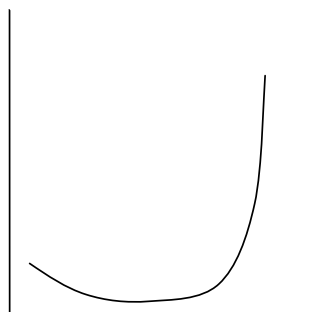
(Кривая унимодального
распределения)



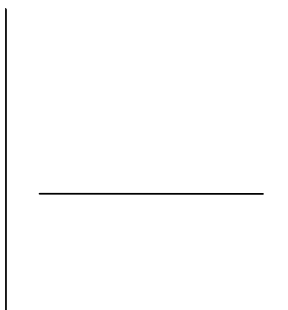
(Кривая бимодального
распределения)



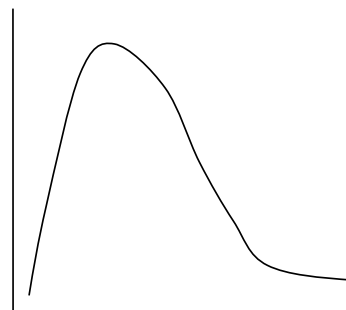
U-образная кривая



J-образная кривая



(Изображение линейного
распределения)



(Влево скошенное
распределение)

Рис. 2.10. Графики частот, позволяющие судить о дисперсии

можно получить другую величину диапазона, другой размах вариации. Увеличивая количество наблюдений, можно либо расширить этот размах, либо оставить его неизменным, однако сократить – нельзя. По этой причине два или более диапазона сравнимы только в том случае, если они состоят из приблизительно одинакового числа наблюдений. Например, не следует сравнивать диапазоны оценок двух студентов, если каждый диапазон не содержит приблизительно одинаковое число этих оценок.

Для дискретных данных процедура измерения диапазона является точно такой же, как и для непрерывных, за исключением того, что истинные пределы становятся при этом в силу необходимости фиктивными. Так, распределение размеров семей от 2 до 12 человек имеет размах вариации равный 11, что является разностью между 12,5 и 1,5. Это значит, что переменная может принимать 11 и только 11 последовательных значений: 2, 3, 4, 5, 6, 7, 8, 9, 10, 11 или 12. В некоторых учебниках эта величина обозначается как «включающий» диапазон и находится по формуле:

$$R = (x_{\max} - x_{\min}) + 1 = 11.$$

Применив эту формулу в случае рассматриваемого выше примера, получим: $(12 - 2) + 1 = 11$.

Значение диапазона определяется весьма просто, подобно любой другой статистической величине, оно даст лишь ограниченную информацию. Поскольку величина его определяется расположением на шкале распределения, необходимо указывать не только его абсолютную величину, но и граничные точки. Повседневная практика подтверждает справедливость этого положения в таких утверждениях, как: «Цена на новые автомобили установится в диапазоне от 4000 до 5000 долларов» или «Температура завтра будет в диапазоне от 12° С до 18° С. Следует обратить внимание, что, например, перечень окладов от 500 до 1000 гривен имеет тот же абсолютный размах вариации, что и перечень от 2500 до 3000 гривен, однако будет с совершенно иным содержанием. А при выборе места для отдыха недоста-

точно знать, что диапазон температур там равен 30°C , важно знать абсолютные значения крайних температур.

Характерной особенностью диапазона является то, что он не учитывает структуры вариации в своих пределах, однако иногда именно эта структура представляет наибольший интерес. Диапазон среднегодового дохода семьи в Украине составляет величину свыше 3500 гривен, что не дает никакого ответа на вопрос, сгруппированы ли доходы в середине, концентрируются ли они к концу интервала, ограничивающего диапазон, или же они равномерно распределены по всему диапазону.

Более того, в большинстве наблюдаемых распределений пределы вариации соответствуют весьма малым частотам; следовательно, полный диапазон, определяемый с помощью этих значений, может создать впечатление большей величины вариации, нежели в действительности. Устанавливая возрастной диапазон студентов вуза от 14-летнего вундеркинда до 64-летней бабушки, которая желает присутствовать в вузе вместе со своими внуками, получаем диапазон в 51 год. Но этот результат затмевает тот факт, что большинство студентов отличаются друг от друга лишь несколькими годами, и, следовательно, полный диапазон, как показатель вариации, в данном случае вводит в заблуждение.

Промежуточные диапазоны. Такая зависимость от пределов вариации может быть уменьшена с помощью промежуточных диапазонов, не учитывающих крайние значения переменной. Ограничивая область наибольшей концентрации наблюдений, такой диапазон обеспечивает большую стабильность и надежность. В обычной практике берется разность между 90-м и 10-м центилями, то есть устанавливается диапазон, включающий 80% случаев. Еще более ограниченным диапазоном является промежуточный интервал между первым и третьим квартилями, который содержит в среднем 50% случаев. Обычно он называется интерквартильным диапазоном.

Интерквартильный диапазон может быть изображен графически путем нанесения квартилей на базовую линию графика. Подобная процедура обнаруживает степень «сгруппированности» случаев, относящихся к среднему интервалу, вокруг

медианы. В рассматриваемом примере четыре интервала, образуемые квантилями, очень сильно различаются по ширине. Хотя каждый из четырех интервалов содержит ровно до 25% общей частоты. Такая классификация наблюдений обратна по отношению к ортодоксальной таблице частот, в которой интервалы выбираются равными, а частоты не равны. Здесь же устанавливаем равные интервалы частот, но допускаем переменную ширину интервала.

Нетрудно убедиться, что построение промежуточных диапазонов, таких, как интервал 10–90% или интерквартильный диапазон, не ограничено никакими рамками. Промежуточные диапазоны в ряде случаев обеспечивают большую ясность в вопросе об относительной концентрации или дисперсии случаев по сравнению с полным диапазоном. Во многих случаях использование стратегически расположенных квантилей (это общий термин мер расчленения) вполне удовлетворительно описывает дисперсию, делает ненужными более усложненные методы.

Отклонение от среднего как мера вариации. Диапазон полный, или промежуточный, позволяет судить об области распределения переменных. Хотя эта мера и имеет некоторое применение, особенно, когда указываются абсолютные границы, она дает информацию лишь о пределах вариации, а не о вариации в этих пределах. Следовательно, можно еще раз повторить, что размах вариации измеряет границы рассеяния, а не полную величину вариаций совокупности. Необходимы некоторые показатели, которые отражали бы степень вариации величин полного распределения.

Попытаемся определить вариацию как отклонение от некоторого значения. Можно, например, получить множество парных разностей для всех величин распределения, а затем «конденсировать» эти разности в некоторый показатель вариации. Однако каждый, кого интересует вариация ряда значений, интуитивно рассматривает их в связи с фиксированным стандартом, который построен на совокупном социальном опыте. «Высокий оклад», «чрезвычайно высокий уровень рождаемости» или «слабые

способности» – все эти суждения являются оценкой относительной величины вариации, замеряемой от некоторой установленной основы. Высокий темп рождаемости можно рассматривать как отклонение в положительном направлении, низкий темп рождаемости – как отклонение в отрицательном направлении от нормы. Следовательно, эта норма становится центральной величиной, относительно которой воспринимается и измеряется вариация.

Точно так же можно судить об уровне смертности, об учебных оценках или о преподавательских окладах, особенно, когда они достаточно близки к наблюдаемым экстремумам. Тем не менее более обоснованно выбирать в качестве точки отсчета не субъективные «стандарты», а средние величины. Другими словами, исследователи стремятся организовать наблюдения вокруг среднего значения, взятого в качестве нормы, полагая, что это среднее является характерным значением. Остается решить следующие проблемы: 1) от какого среднего вычислять отклонения; 2) как представить эти отклонения компактным показателем.

Выбор нормы (от которой «отклоняется» отклонение). Вообще, нормой в данном случае считается средний показатель. Однако нам известны, по крайней мере, три разные средние меры ($M[X]$, Mo , Me), а выбрать в качестве нормы нужно только одну из них. Очевидно, что для симметричных унимодальных распределений этот выбор не представляет особых трудностей, поскольку среднее арифметическое, мода и медиана оказываются равными. Однако совершенно симметричные распределения встречаются крайне редко. Следовательно, дилемма возникает при выборе нормы для тех распределений, для которых мода, медиана и среднее арифметическое отличаются друг от друга. Многие исследователи выбирают моду, то есть максимальную частоту в качестве основы для сравнения. Однако более широко в этих целях используются среднее арифметическое или медиана, которые считаются более репрезентативными.

Как уже говорилось, среднее арифметическое – это величина, вариации которой в обе стороны равны. Следовательно, среднее удобно использовать в качестве нормы, однако и медиана столь

же успешно может служить началом отсчета вариации. Поскольку медиана делит совокупность распределения на равные части – сумма отклонений от медианы меньше, чем от любой другой точки. Говоря другими словами, медиана – это точка, для которой арифметические «ошибки» минимальны. Приведенное утверждение можно назвать принципом минимального отклонения.

Абсолютные показатели вариации. Среднее линейное отклонение ($\bar{d}$). Простая сумма арифметических отклонений является бесполезной в качестве показателя вариации, поскольку она непосредственно зависит от количества объектов в распределении. Например, она будет большой для 1000 объектов и малой для 10. Чтобы устранить этот фактор, разделим сумму отклонений на N и, следовательно, будем измерять отклонение, приходящееся на один объект. Этот результат называется средним линейным отклонением ($\bar{d}$) и вычисляется с использованием формул:

$$1. \text{ Для первичного ряда: } \bar{d} = \frac{\sum (x_i - M[X])}{n}.$$

$$2. \text{ Для вариационного ряда: } \bar{d} = \frac{\sum |x_i - M[X]| \times n_i}{N}.$$

$$3. \text{ Для интервального ряда: } \bar{d} = \frac{\sum |x_{ci} - M[X]| \times n_i}{N}.$$

При этом

- x_i – значение переменной (ее величина);
- x_{ci} – средняя точка интервала (в случае интервального ряда);
- $M[X]$ – среднее арифметическое (причем вместо $M[X]$ мы можем подставить в формулу Me , в том случае, если распределение неравномерно, асимметрично);
- n_i – частота значения переменной (или частота интервала в случае интервального ряда);
- n – число вариантов признака (в случае первичного ряда);
- N – сумма всех частот.

Квадратичные отклонения. Измерение вариации с помощью простых арифметических (линейных) отклонений от центрального значения является наиболее простой процедурой. Если исследователя интересует лишь наличие или отсутствие дисперсии, то легко вычисляемый d_{cp} вполне подходит для этой цели. Однако вариация, как правило, измеряется через квадратичные отклонения от средних величин. Логика этой операции основывается на принципе минимальных отклонений. Как сумма квадратичных отклонений от среднего арифметического, так и сумма квадратичных отклонений от медианы является минимальной величиной. Это положение получило название «принцип наименьших квадратов» и является одним из наиболее важных принципов статистических расчетов.

Метод квадратичных отклонений может показаться, на первый взгляд, искусственным и излишне усложненным. Если вариация может быть удовлетворительно измерена посредством вычисления линейного отклонения, то какое дополнительное преимущество дает возведение в квадрат? Удовлетворительный ответ на этот вопрос можно получить, лишь углубившись в изучение математической статистики. В рамках этого учебника целесообразность такого углубления представляется сомнительной, потому будем полагаться на практический опыт социологов, которые прибегают к вычислению квадратичного отклонения намного чаще, чем простого среднего отклонения

Квадратичные отклонения могут быть получены несколькими способами, каждый из которых пригоден для определенной цели: 1) *сумма квадратов отклонений (дисперсия)*; 2) *вариация*; 3) *среднее квадратическое отклонение*.

Сумма квадратов отклонений (дисперсия). Дисперсия (δ^2 или D) – величина равная среднему значению отклонений отдельных значений признака от среднего значения. Для того, чтобы найти эту величину, необходимо произвести вычисления по следующим формулам:

$$1. \text{ Для первичного ряда: } \delta^2 = \frac{\sum (x_i - M[X])^2}{n}.$$

$$2. \text{ Для вариационного ряда: } \delta^2 = \frac{\sum (x_i - M[X])^2 \times n_i}{N}.$$

$$3. \text{ Для интервального ряда: } \delta^2 = \frac{\sum (x_{ci} - M[X])^2 \times n_i}{N}.$$

При этом

- x_i – значение переменной (ее величина);
- x_{ci} – средняя точка интервала (в случае интервального ряда);
- $M[X]$ – среднее арифметическое;
- n – число вариант признака;
- n_i – частота значения переменной (или частота интервала в случае интервального ряда);
- N – сумма всех частот.

В качестве демонстрации процедуры вычисления дисперсии предлагаем обратиться к примеру, который рассматривался в предыдущем параграфе: «Сеть дошкольных учреждений всех ведомств в г. Харькове в 1999 г.». Мы вычисляли среднее арифметическое для сгруппированных данных, пытаясь таким путем найти среднюю численность детей в харьковских дошкольных учреждениях. По результатам этих вычислений мы получили $M[X] = 157,07$. Этот результат и процедура его нахождения представлены в третьем столбце расчетной (для дисперсии) таблицы, представленной ниже (см. Табл. 2.10). Когда среднее нам известно, осуществление последующих действий для нахождения дисперсии не составит большого труда, что также видно из данной таблицы.

Следуя формуле вычисления дисперсии, нам осталось выполнить последнее действие: итоговое табличное значение 243408,78 разделить на сумму всех частот, равную 212. Получаем: $\sigma^2 = 243408,78 / 212 = 1148,15$.

Среднее квадратическое отклонение (сигма – σ). Поскольку вариация основана на квадратических отклонениях, она не является линейной мерой. Если требуется линейная мера, то необходимо извлечь квадратный корень из обеих частей

Таблица 2.10

**Сеть дошкольных учреждений всех ведомств
в г. Харькове в 1999 г.¹⁴**

Количество д/у в 1999 г. (n_i)	В них детей в среднем ($\bar{x}_i$)	$n_i \times \bar{x}_i$	$(x_i - M[X])^2$	$(x_i - M[X])^2 \times n_i$
33	161,06	33? 161,06= =5315	161,06 – <u>157,07</u> = =15,92	525,52
13	140,62	1828	270,75	3519,81
33	141,88	4682	230,77	7615,51
17	182,82	3108	663,24	11275,15
23	93,09	2141	4093,83	94158,09
35	208,46	7296	2640,64	92422,35
24	151,88	3645	26,99	647,71
18	126,00	2268	965,34	17376,21
16	188,56	3017	991,78	15868,44
N=212		33300/212 = <u>157,07</u> ($M[X]$)		243408,78

соотношения. В такой форме эта величина известна как среднее квадратическое отклонение или сигма (σ). Среднее квадратическое отклонение показывает, насколько в среднем каждое значение признака отклоняется от среднего. Геометрически, при нанесении на график, сигма показывает, насколько кривая распределения «размыта» относительно среднего.

Сигма используется как мера вариации, совершенно аналогично $\bar{d}$, от которого отличается главным образом тем, что отклонения квадратичны, а их среднее – линейно. Однако извлечение корня не может полностью уничтожить влияние предшествующего возведения в квадрат; эффект взвешивания частично сохраняется.

¹⁴ Стат. збірник: показники роботи закл. освіти та наук. установ обл. за 1999 рік. [за заг. редакцією О. Л. Сидоренка, А. С. Доценка, П. С. Деметсьва]. – Х., 2000. – 82 с.

Вычисление среднего квадратического отклонения (σ).
В принципе среднее квадратическое отклонение является несколько более сложным показателем, нежели среднее линейное отклонение, требуя дополнительного возведения отклонений в квадрат и извлечения квадратного корня из их среднего. Вообще, если нам известна дисперсия, то сигма находится очень просто, путем извлечения квадратного корня, а именно:

$$\sigma = \sqrt{\delta^2}.$$

Если же дисперсия распределения неизвестна, то среднее квадратическое отклонения для этого распределения находится по следующим формулам:

$$1. \text{ Для первичного ряда: } \sigma = \sqrt{\frac{\sum (x_i - M[X])^2}{n}}.$$

$$2. \text{ Для вариационного ряда: } \sigma = \sqrt{\frac{\sum (x_i - M[X])^2 \times n_i}{N}}.$$

$$3. \text{ Для интервального ряда: } \sigma = \sqrt{\frac{\sum (x_{ci} - M[X])^2 \times n_i}{N}}.$$

При этом

- x_i – значение переменной (ее величина);
- x_{ci} – средняя точка интервала (в случае интервального ряда)
- $M[X]$ – среднее арифметическое;
- n – число вариант признака;
- n_i – частота значения переменной (или частота интервала в случае интервального ряда);
- N – сумма всех частот.

«Среднее линейное VS среднее квадратическое». Для описательных целей прием отбрасывания знаков при вычислении $\bar{d}$ является совершенно законным. Поскольку $\bar{d}$ измеряет отклонения без возведения в квадрат, то оно по абсолютной величине меньше, нежели σ , которое непропорционально увеличивает большие отклонения в результате возведения в квадрат.

Однако пренебрежение знаками делает $\bar{d}$ непригодным для использования в последующих алгебраических вычислениях, независимо от его начала отсчета. Именно поэтому одним из наиболее распространенных средств статистического анализа в течение почти столетия была σ , причем не только как мера дисперсии, но и как составная часть более сложных вычислений. Широкое использование сигмы в какой-то мере объясняется двумя присущими ей достоинствами. *Во-первых*, σ дважды отражает величину каждой переменной распределения: а) точка отсчета, от которой измеряются отклонения ($M[X]$), сама является репрезентацией всех переменных; б) каждая величина как таковая представлена квадратичным отклонением. *Во-вторых*, возведение отклонений в квадрат автоматически снижает проблему знака отклонения.

Относительные показатели вариации.

Линейный коэффициент вариации. Как уже говорилось, любое отклонение имеет смысл только лишь тогда, когда известно, от чего оно будет «отклоняться», то есть лишь после того, как будет задано начало отсчета или норма. Этот принцип воплощен в коэффициенте вариации, который выражает меру вариации через процентное отклонение от начала отсчета, независимо от того, является ли оно медианой или средним арифметическим. В том случае, если $\bar{d}$ отсчитывался от среднего арифметического, формула коэффициента вариации в знаменателе будет иметь $M[X]$, если же он основан на медиане, формула коэффициента вариации в знаменателе будет иметь Me .

$$V_d = \frac{\bar{d}}{M[X]} \times 100$$

(если мы хотим отобразить коэффициент в %%).

При этом

- $\bar{d}$ – среднее линейное отклонение;
- $M[X]$ – среднее арифметическое (в делителе может быть и Me , в случае неравномерного, асимметричного распределения).

Подобно среднее квадратическое отклонение также может

быть превращено в меру относительной вариации посредством нормирования его по отношению к собственному началу отсчета, то есть среднему арифметическому:

$$V_{\sigma} = \frac{\sigma}{M[X]} \times 100$$

(если мы хотим отобразить коэффициент в %%).

При этом

- σ – среднее квадратическое отклонение;
- $M[X]$ – среднее арифметическое для данного распределения.

Коэффициент вариации является показателем изменчивости признака относительно его средней величины. К примеру, в результате соответствующих вычислений, мы получили $V_{\sigma} = 0,8$. Если выразить его в процентах – имеем 80%. А это значит, что только 20% всего распределения по данному признаку относительно однородно и приближено к среднему значению. Остальная же часть распределения неоднородна, 80% всех значений очень сильно отличаются от среднего и «далеко» рассеяны по отношению к этому среднему.

Коэффициенты вариации оказываются особенно полезными в процедурах сравнения, поскольку они не зависят от абсолютных значений и от употребляемых единиц измерения. Коэффициент вариации позволяет сравнивать множества малых и больших однородных величин, а также до определенной степени и качественно отличных объектов. Однако *V применим только в тех случаях, когда: (1) наблюдаемые значения имеют нуль; (2) все интервалы равны.* Кроме того, *V* более целесообразно использовать при сравнениях между последовательностями связанных данных. Относительная вариация заработной платы на Востоке Украины может быть меньше, чем на Западе Украины. При измерении симпатии публики к некоторому композитору высокий *V* будет получаться при большом расхождении в мнениях; низкий коэффициент, наоборот, отражал бы тенденцию согласия.

Сопоставление мер средней тенденции и вариации, интерпретация результатов такого сопоставления. Следует

подчеркнуть, что малое значение σ при большом среднем указывает на большую однородность данных и в силу этого на типичность среднего, что в некоторых условиях крайне существенно. Среднее, равное 125, при $\sigma=5$ и V , равным 4%, более репрезентативно, нежели среднее в 125, при $\sigma=25$ и V равным 20%. V , равное нулю, указывает на отсутствие вариации вообще. Следует отметить, что в той мере, в какой σ увеличивает вариацию относительно среднего арифметического, V соответственно увеличивает относительную вариацию.

2.7. Вариация качественных переменных

Вариация качественных переменных. Очевидно, что для номинальных признаков некорректным является использование всех приведенных выше мер разброса. У качественных переменных не существует «нуля» отсчета, и, следовательно, они не имеют величин. Не существует среднего значения диапазона и промежуточных интервалов. Следовательно, не существует и арифметических отклонений.

Это, однако, не означает, что любая группа качественных переменных состоит из совершенно идентичных событий. Попытаемся понять, как можно интерпретировать такой разброс. Два события можно считать различными, если они не обладают различными качествами. Вместо вычисления величин, подсчитываются различия в качествах. Чем больше число различных пар событий, тем более неоднородна совокупность, и, следовательно, тем больше вариация внутри нее. Аналогично, чем меньше это число, тем больше однородность внутри совокупности и меньше вариация. Поэтому разумно установить показатель качественной вариации по полному числу различных пар событий данного множества. Вопрос теперь лишь в том, *во-первых*, как подсчитать полное число различий и, *во-вторых*, как превратить это число в компактный показатель.

Чтобы найти полное число различий, суммируются всевозможные различия в группе событий. Например, в множестве из

шести мальчиков и шести девочек каждый из шести мальчиков будет отличаться по своим признакам от каждой из шести девочек, давая в итоге 36 различий полов. Если бы имелось девять мальчиков и три девочки, то каждый из девяти мальчиков отличался бы от каждой из трех девочек, что давало бы в итоге 27 различий. В группе 12 мальчиков очевидный результат отсутствия различий был бы получен при умножении 12 на ноль.

Очевидно, что процедура определения полного числа различий сводится к следующему правилу: умножаем частоту каждого признака на частоту каждого отличного от него признака и суммируем эти произведения. Например, в совокупности из четырех католиков, пяти христиан и шести иудеев будем иметь: $(4 \times 5) + (4 \times 6) + (5 \times 6) = 74$ различия.

Коэффициент качественной вариации (V_q). Число различий, как показатель вариации, сравнимо только с максимально возможным числом различий. Это максимальное число различий будет наблюдаться в том случае, когда все частоты различных признаков равны. Таким образом, максимум вычисляется путем приравнивания частот (т. е. вычисления средней частоты), перемножения частот и суммирования произведений. Другими словами, осуществляются следующие операции: (1) находится средняя частота; (2) этот результат возводится в квадрат; (3) квадрат умножается на число возможных пар признаков. В вышеупомянутом примере девяти мальчиков и трех девочек максимально возможное число различий пола в группе из 12 было бы: $6 \text{ (мальчиков)} \times 6 \text{ (девочек)} = 36$, или в этом конкретном случае средняя частота умножается на саму себя.

Относительная величина вариации теперь может быть измерена с помощью отношения между наблюдаемым числом различий его гипотетическим максимумом:

$$\text{Коэффициент качественной вариации } (V_q) = \frac{\text{Полное число набл. различий}}{\text{Макс. возможн. число различий}}.$$

Результат вычислений по соответствующей формуле представляет собой долю всех значений признака, которые сильно неоднородны. Очевидно, что этот показатель принимает значения

от 0 до 1. Чем больше коэффициент вариации приближен к нулю, тем меньше вариация значений признака. Как и в случае с вариацией количественных данных, полученное число, которое не может превышать единицы и опускаться ниже нуля, можно представить в %%, умножив полученный результат на 100.

Проиллюстрируем применение данной формулы на предыдущем примере с 9 мальчиками и 3 девочками:

$$V_q = \frac{27}{36} = 0,75(\times 100) = 75\%.$$

Этот показатель говорит о том, что вариация значений признака довольно высока, и только 25% (100% минус 75%) относительно однородны.

Среднее число членов каждой из трех упомянутых выше религиозных групп равно пяти. Умножив «5» на «5» и просуммировав эти три произведения, найдем, что максимальное число различий должно быть равно 75. Наблюдаемые различия, как уже было вычислено, равны 75. Следовательно:

$$V_q = \frac{74}{75} = 0,99(\times 100) = 99\%.$$

Полученная цифра говорит о предельно высокой неоднородности всех значений признака.

Примеры использования коэффициента вариации. Как уже говорилось, рассматриваемый нами коэффициент может быть использован для сравнения тех или иных относительных величин. Например, попытаемся сравнить, насколько вырос/упал уровень «остепененности» системы высшего образования Украины с 1996 по 2000 год (см. *Табл. 2.11*).

Таблица 2.11

Динамика численности основного профессорско-преподавательского состава высших учебных заведений Украины III–IV уровня аккредитации

Годы		Профессора	Доценты
1995	1996	29597	5728
1999	2000	28540	6546

В Украине в 1995/96 учебном году в высших учебных заведениях III–IV уровня аккредитации работало 29597 профессоров и 5728 доцентов, следовательно, максимально возможное число различий будет равно $((29597+5728)/2)^2=17663^2$, тогда

$$V_d = \frac{29597 \cdot 5728}{17663^2} = 0,54(\times 100) = 54\%.$$

В 1999/2000 учебном году ситуация была такова:

$$V_q = \frac{28540 \cdot 6546}{17543^2} = 0,61(\times 100) = 61\%.$$

Элементарное нормирование. Необходимость нормирования. Любое событие исследователь рассматривает не изолированно, а в сравнении с конкретной нормой, вытекающей из социальной основы данного события. Например, годовой доход в 3500 гривен воспринимается социологом не как отвлеченное число, а как социальное явление, отнесенное к стандарту, основанному на опыте исследователя; факт рождения ста человек в общности имеет смысл лишь в связи с такими данными, как количество населения, период времени, число рождений в предшествующий год или число рождений в других общностях. Если норма, с которой производится сравнение, не установлена точно, то исследователь невольно установит ее самостоятельно.

Еще большие трудности возникают при сравнении двух и более величин, взятых из различных совокупностей. Например, сравнение умственных способностей мужчин и женщин, проводимое на основе тестов, может быть ошибочным, поскольку известно, что результаты таких тестов зависят от уровня образования, который по полу может быть распределен неравномерно.

Одна из важнейших функций математической статистики по отношению к социологии заключается в предоставлении метода, позволяющего обоснованно сравнивать несколько величин. Существует много способов решения этой задачи; некоторые из них будут рассматриваться в данном подразделе. Совокупность этих процедур можно назвать операциями нормировки,

поскольку они устанавливают определенные стандарты наблюдаемых величин. Процесс нормировки уже известен читателю из опыта расчета вариаций, и в настоящей главе остановимся на некоторых других аспектах этой важнейшей статистической процедуры, особенно в применении к социологическим материалам.

Можно осуществлять нормировку приблизительно в следующем порядке сложности: 1) *процентные отношения*; 2) *пропорции*; 3) *степени*; 4) *индексы*; 5) *подклассификация*; 6) *стандартизация*.

Процентные отношения. Простейшая форма нормировки состоит в приведении рядов абсолютных чисел к стандартной численной основе. Громоздкие абсолютные числа заменяются указанием их отношения к некоторой основе, выраженной в процентах. Вместо того чтобы указывать, что зарегистрированное число юношей и девушек – студентов вуза равно соответственно 9244 и 4622, превращают эти значения в 66,7 и 33,3%. Эта операция настолько привычна, что основной принцип, на котором она основана, не всегда полностью осознается.

Пропорции. Можно сравнивать две величины в форме отношения или выражать одну из них как кратное другой. Отношения бывают различными по составу: можно выделить отношения «часть – часть» частот в пределах одного и того же множества и отношение «целое к целому» между частотами двух взятых переменных. Таким образом, соотношение полов может рассматриваться как отношение «часть – часть», оно сравнивает число мужчин в данной совокупности с числом женщин.

В 1999 году в Украине было 34 млн человек городского населения и 16,1 млн человек – сельского. Соотношение между этими громоздкими числами легче запомнить, если выразить его в следующем виде: 2,1 городского жителя приходится на 1 сельского.

В антропометрии цефалический индекс представляет собой отношение «целого к целому» двух черепных мер – ширины к длине – с целью количественного различия между круглой и овальной формами головы. Отношение умножается на 100, чтобы сделать запись более ясной:

$$\text{Цефалический индекс} = \frac{\text{ширина}}{\text{длина}} \times 100\%.$$

Аналогично этому степень умственного развития равна отношению между умственным и хронологическим возрастом испытуемых, что позволяет сравнивать людей различных возрастов по умственному развитию. Таким образом:

$$Iq = \frac{\text{Умственный возраст}}{\text{Хронологический возраст}} \times 100\%.$$

Это отношение равно 100 в том случае, когда хронологический возраст и умственный оказываются равными. Другие общепринятые пропорции, используемые в социально-экономическом анализе, — это пропорции: люди — жилье; население — земля; дети — взрослые.

Степени (коэффициенты). Степень является по существу арифметическим средним. Она представляет собой среднее число значений одной переменной, выраженное в единицах другой. Так, степень, равная 20 километрам на 1 литр, есть среднее потребление топлива, при котором расстояние в километрах будет принимать определенное значение при замене одного литра другим. Хотя все степени основаны на прошлых наблюдениях, они обеспечивают предсказание будущего. Поэтому иногда вычисленные степени и средние могут рассматриваться как ожидаемые величины. Будучи в основном результатом большого числа наблюдений, степень часто оказывается эффективным инструментом социологического анализа. Многие степени стали общепринятыми понятиями в области социологии: степень семейности, степень преступности, степени рождаемости и смертности и многие другие видоизменения этих степеней, которые являются рабочими понятиями в социологии.

Статистический смысл степени зависит в основном от двух переменных: от проблемной переменной и от нормированной переменной. При вычислении степеней наиболее важным является выбор нормированной переменной, с которой будет сравниваться проблемная переменная. Например, при вычислении

степени рождаемости необходимо выбрать нормирующую совокупность, с которой в дальнейшем будет сравниваться абсолютное число актов рождения. Для этой цели можно было бы использовать либо полную совокупность людей, либо число женщин, достигших зрелого возраста, либо число замужних женщин, достигших зрелого возраста. Наиболее часто используется, хотя и не вполне оправданно, полная совокупность людей: мужчин, женщин, детей.

Следующим этапом построения степени является выбор стандартной численной основы: 10, 100, 1000 или кратное им значение. Назначение числовой основы состоит просто в указании десятичного масштаба для удобства табулирования, для облегчения цитирования и более быстрого понимания. Численный масштаб часто устанавливается по соглашению, особенно когда не существует другого выхода, кроме простого подчинения установившейся традиции. Так, степень рождаемости, равная 24, имеет интернациональный смысл и означает 24 случая рождений на 1000 человек всего населения в данном году и на данной территории. Такое представление воспринимается более ощутимо, нежели запись: 768 из 32462. Вычисление осуществляется следующим образом: число рождений = 768, полное население = 32462, числовой масштаб = 1000,

$$\text{Степень рождаемости} = \frac{768}{32462} \times 1000 = 24.$$

Обобщенная формула читалась бы следующим образом:

$$\text{Степень} = \frac{\text{Частота проблемной переменной}}{\text{Частота нормировочной переменной}} \times \text{Числовой масштаб.}$$

В обозначениях:

$$\text{Степень} = \frac{r_{nn}}{r_{mn}} \times r_m,$$

- где • r_{nn} — частота задачи;
 • r_{mn} — нормирующая частота;
 • r_m — числовой масштаб.

Ни нормирующая переменная, ни числовой масштаб не определены твердым соглашением, как, например, для степени рождаемости или смертности. В таких случаях допускается некоторая свобода действий, однако результаты должны снабдиться примечаниями. Степень разводов может быть вычислена как для всего населения, так и для числа супружеских пар в один и тот же год и на одной площади, или даже для числа браков в течение предшествующих десяти лет, как это иногда делается для большинства разводов. Степени преступности могут вычисляться для определенного возраста и конкретных половых групп; степень браков – лишь для возрастной группы от 14 лет и более.

Из предыдущих примеров можно сделать вывод о том, что нормированная переменная должна иметь те же характеристики, что и проблемная переменная – вместе они объединяются общим названием «открытая группа». Если смерть может случиться с каждым человеком, то рождение ребенка, женитьба или замужество и развод – не с каждым. Следовательно, степень, основанная на разумно выбранной открытой совокупности, менее подвержена искажениям из-за внешних факторов. Степени рождаемости, браков и разводов, вычисленные по отношению к полному населению, обычно называются приближенными степенями, а степени, вычисленные для особых групп, обозначаются как «удельные» или «уточненные». Преимущество приближенной степени состоит в ее простоте и удобстве для ориентировочных расчетов. Уточненные степени необходимы для профессионального социологического исследования.

Индексы. В статистике индекс – термин, используемый в разговорном языке и технике для самых различных типов мер, – обычно относится к более сложным степеням или множествам пропорций. В качестве вторичной меры он обычно предназначается для описания вариации, которая в непосредственном виде могла бы быть совершенно незаметной. В более формализованном варианте он обычно описывает отношение между двумя величинами, одна из которых взята в качестве нормы, или ожидаемой величины, тогда как другая является измеряемой величиной.

Так, индекс стоимости жизни сравнивает цены в конкретном году со средними ценами для «нормального» года. Индекс, равный 139 в 1995 году, при использовании в качестве основного года – 1991 показывает, что стоимость жизни повысилась на 30% по сравнению с основным годом, для которого стоимость принята равной 100. Хотя такой обманчиво простой индекс может бойко цитироваться любым журналистом, его внутреннее содержание, включающее охват, взвешивание, метод усреднения наблюдений, а также выбор основного периода, свидетельствует о его статистической сложности. Аналогично этому V_q сравнивает наблюдаемое число различий признаков заданного множества с максимально возможным числом различий, которое в данном случае служит в качестве нормы.

Таблица 2.12

Вычисление индекса социально-экономической классификации, студентов университета и населения региона

Группа	Университет		Регион %	Разность	Индекс
	кол-во	%%			
Техники	408	19,3	4,7	14,6	411
Инженеры, управленцы	507	24,0	4,7	19,3	511
Бизнесмены	290	13,7	5,0	8,7	274
Клерки	286	13,5	12,8	0,7	105
Фермеры	196	9,3	15,9	-6,6	58
Квалифицированные рабочие	267	12,6	17,0	-4,4	74
Полуквалифицированные рабочие	94	4,5	19,9	-15,4	23
Неквалифицированные рабочие	65	3,1	20,0	-16,9	15
<i>ИТОГО</i>	2,113	100%	100%		

Таблица 2.12 показывает построение и использование индекса для измерения социальной стратификации студентов университета, а также то, насколько представлены в нем различные социальные классы. Логическая основа построенных в этой таблице индексов следующая: дочери состоятельных родителей составляют 19,3% от общего числа девушек в университете, тогда

как в другом регионе самая состоятельная группа составляет 4,7%. Соотношение между такими парами процентных отношений могло бы служить мерой доступности обучения для различных классов. Существует два метода, с помощью которых можно было бы измерять эти различия:

- 1) простого различия между процентными отношениями;
- 2) нормированных индексов.

Что касается первой альтернативы, то анализ различий обнаруживает, что при перемещении по социальной шкале различия возрастают. Однако эти различия являются абсолютными, в силу чего их невозможно нормировать по величине или началу отсчета. Чтобы нормировать данные различия, строится индекс, что составляет содержание второй альтернативы.

Таким образом, если бы посещаемость университета распределялась случайным образом между всеми классами населения то можно было бы ожидать, что состоятельная группа, которая составляет 4,7% всего населения, должна дать вклад, равный 4,7% от всего числа всех студентов. Ожидаемая и наблюдаемая доля студентов из состоятельных семей были бы в данном случае идентичными, соотношение между ними было бы равно единице. В действительности, состоятельная часть студентов университета

равна 19,3%, что равно $\frac{19,3}{4,7} \cdot 100 = 41,1\%$ соответствующего

процента для региона, или в 4,11 раза больше ожидаемого значения. Таким же образом можно нормировать все другие социально-экономические процентные отношения.

Другой подход к определению индекса заключается в нормировке последовательности величин относительно их среднего. В этом случае индекс представляет собой просто отношение между данной величиной и средним значением последовательности. Например, средняя степень смертности в ряде городов равна 10,5. Если все факторы, влияющие на степень смертности, были бы одинаковыми во всех городах, тогда все степени смертности станут идентичными и, следовательно, будут равны среднему значению последовательности. Поскольку при данном

предположении среднее является ожидаемым или теоретическим значением, то, чтобы оценить вес факторов, влияющих на различия, индивидуальные степени измеряются по отношению к среднему. Например, если наблюдаемая степень смертности для города равна 7, то мы могли бы вычислить индекс следующим образом:

$$\text{ИНДЕКС} = \frac{\text{наблюдаемая степень}}{\text{ожидаемая степень}} \times 100\% = \frac{7}{10,5} \times 100\% = 67\%.$$

Это значит, что степень смертности в данном городе составляет 67% от средней. С помощью этого метода любой город можно расположить на шкале отношений к среднему. Этот прием нормировки аналогичен V , поскольку он выражает исходные значения через их собственное среднее.

По аналогии находится и, так называемый, сезонный индекс смертности населения, который вычисляется, как отношение между месячной мерой и среднегодовой мерой. Этот индекс используется для измерения флуктуаций рождаемости, смертности, промышленной продукции и некоторых других экономических показателей.

Нормировка посредством подклассификации. Подобно тому, как отдельное абсолютное значение не имеет смысла до сопоставления с соответствующей нормой, точно так же и степень или процентное отношение практически не имеет социологического смысла, пока они рассматриваются самостоятельно. Они приобретают смысл тогда, когда их сопоставляют с аналогичными степенями, например, сравнивая степень рождаемости двух регионов и степень бракосочетаний католиков и христиан. Однако не следует делать слишком поспешного вывода о существовании причинно-следственного соотношения между такими спаренными переменными. Наблюдаемые вариации степеней могут иногда возникать в результате действия факторов, не учтенных в классификации. Такие факторы можно назвать скрытыми. Во многих случаях сравнение двух или большего числа наблюдаемых величин искажаются в результате воздействия именно таких скрытых факторов. Так, более высокая степень рождае-

мости в одном регионе может произойти не в силу большей плодовитости его населения, а из-за большего процентного количества женщин, способных к деторождению. Этот случайный фактор не классифицирован в приближенной степени.

Под нормировкой посредством подклассификации обычно подразумевается разделение факторов на «внешние» и «внутренние», причем внешние факторы не должны изменяться в ходе исследования. Из всего сказанного следует, что нельзя распространять на все группы результаты, полученные для соответствующих подгрупп, ведь они отличаются не только весом, но и другими факторами. Поэтому необходимо разработать метод, который позволил бы получить простую, уточненную, но свободную от влияния весов, степень. Соответствующий метод получил название «стандартизация», а получаемая степень – «стандартизованная».

Было обнаружено, например, что повышенную степень преступности населения можно объяснить более высокой долей молодежи в нем. Именно этим, а не избыточной тенденцией к совершению преступлений, объясняется повышенная степень преступности местного населения. В целом же стандартизованные степени преступности для каждой группы населения строятся следующим образом: вычисляют ожидаемое число преступлений для местного населения при условии, что оно имеет такой же возрастной состав, что и другие поселенческие группы. Другими словами, необходимо действовать так, как будто бы все интересующие нас поселенческие группы, имеют одинаковое распределение по возрасту.

Хотя стандартизация кажется полностью формализованной процедурой, нет никакой формулы, которая предписывала бы, насколько подробной и какой именно должна быть подклассификация. Поэтому социолог не освобождается от содержательного анализа задачи. Например, возраст можно было бы подклассифицировать более чем на три интервала; чем больше число подразделений, тем больше точность сравнений. Однако существуют практические ограничения, за пределы которых распространять подклассификацию нет необходимости. Иногда достаточно

самого грубого разделения возраста на три интервала, чтобы установить важность возрастного фактора в степени преступности.

Применение стандартизации не ограничивается степенями и процентами. Любой вид арифметического среднего может быть стандартизован при наличии необходимых данных для подклассификации. Неограниченные возможности стандартизации напоминают еще раз, насколько далека «окончательная истина» от данных, которые лежат перед нами и на которых, тем не менее, часто основываются мнения и действия.

Довольно часто приближенные степени, которые необходимо превратить в нормированные, соответствуют большим территориям и охватывают большие интервалы времени. Трудно сравнивать статистику разных стран, если отсутствует всеобщее соглашение о стандарте. В качестве такого стандарта в 1901 г. часто применялся английский «стандартный миллион». Так как распределение населения по возрасту – один из наиболее искажающих факторов, в интерпретации социальной статистики «стандартный миллион» – есть возрастное распределение одного миллиона британского населения. Такая процедура нормирует степени по возрасту, тем самым превращая их в величины, удобные для сравнения.

Перекрестные таблицы обычно создаются с целью выявления статистических ассоциаций, однако из-за присутствия скрытых факторов, полученные значения ассоциаций не следует рассматривать как безусловно достоверные. Чтобы выявить скрытые факторы, можно подклассифицировать перекрестные таблицы. Для иллюстрации рассмотрим данные (фиктивные) для 52 районов, перекрестно классифицированных близостью к определенному типу производства и по степени правонарушений (см. *Табл. 2.13*).

Эта таблица указывает, что промышленные районы чаще характеризуются высокими степенями правонарушений, чем районы сельских жителей, указывая, как будто на статистическую связь между местом проживания и преступностью. Эта связь выявляется более четко в процентном распределении,

Таблица 2.13

**Приближенные коэффициенты преступности,
промышленных городских и сельских районов**
(возраст населения от 15 до 75 лет)

<i>Тип района</i>	<i>Коэффициент преступности</i>					
	<i>Количество</i>			<i>Процент</i>		
	<i>Высокий</i>	Низкий	Сумма	<i>Высокий</i>	Низкий	Сумма
<i>Промышленный</i>	15	8	23	65	35	100
<i>Сельский</i>	10	19	29	34	66	100
<i>ИТОГО</i>	25	27	52	48	52	100

которое показывает, что 65% всех промышленных районов находятся в категории «высокой преступности», тогда как подклассифицированные аналогичным образом сельские районы составляет в этой категории только 34%. С точки зрения приближенной степени, вывод о связи между местом проживания и преступностью кажется убедительным. Однако такой вывод неприемлем для любого специалиста по социальной патологии города. Он указал бы, что жители одних районов сосредоточены в регионах с низким уровнем жизни, тогда как жители других районов чаще проживают в более благоприятно расположенных районах с относительно высоким уровнем жизни. Разумно поэтому спросить, будет ли сохраняться разница между городскими и сельскими районами в отношении преступности, если эти районы нормировать на одинаковый экономический уровень. Чтобы ответить на этот вопрос, необходимо подклассифицировать районы согласно жизненным стандартам и провести сравнения в пределах одинаковых социально-экономических подклассов (см. *Табл. 2.14*).

Анализируя эту таблицу, можно увидеть, что связь между местом проживания и преступностью исчезает: из 24 экономически худших районов 21 район или 88% находятся в категории высокой преступности (и это справедливо как для промышленных, так и для сельских районов). Из 28 районов, более развитых экономически, только 4 или 14% находятся в категории

Таблица 2.14

**Коэффициенты преступности,
отнесенные к возрасту (15–75 лет) и уровню жизни респондентов,
проживающих в промышленных и сельских районах**

Тип района	Уровень жизни					
	Высокий уровень			Низкий уровень		
	Коэффициент преступности			Коэффициент преступности		
	Высокий	Низкий	Итого	Высокий	Низкий	Итого
Промышленный	1	6	7	14	2	16
Сельский	3	18	21	7	1	8
Итого:	4	24	28	21	3	24
Процентное распределение						
Промышленный	14	86	100	88	12	100
Сельский	14	86	100	88	12	100
Итого:	14	86	100	88	12	100

высокой преступности независимо от типа района проживания. В результате в этой гипотетической иллюстрации преступность полностью зависит от экономического уровня и совсем не зависит от типа района проживания. Такая связь называется ложной, так как она фактически является результатом действия скрытого социально-экономического фактора, который дает «подлинное» объяснение.

В основном подклассификация – процедура для уточнения сравнений, она, подобно анатомическому расчленению, является инструментом для более полного и глубокого статистического анализа. Подклассификация отличается от стандартизации тем, что она чисто описательна и имеет столько же показателей, сколько выбрано подклассов. Стандартизованная степень, с другой стороны, является скорее гипотетической, чем описательной; это единый, совокупный показатель, взвешенный по ряду частот подклассов, используемых как стандарт.

Вопросы и задания для самоконтроля

1. Дайте определения следующим терминам: первичный, вариационный и динамический ряд, дискретный и интервальный ряд.

2. Какова разница между частотой и частостью?

3. Дайте определения следующим терминам: абсолютная, относительная, накопленная частота, совместная частота.

4. Распределите показатели из приведенного ниже списка по принципу: абсолютные/относительные частоты: 28 кг; 35%; 246 шт.; 0,45; 55 мин; 1/3.

5. В предыдущем месяце на производственном предприятии была осуществлена аттестация рабочего персонала с применением тестовых методик. В итоге 40 человек из общего числа протестированных работников – 800 аттестацию не прошли. Было принято решение, что этой группе работников необходимо пройти месячный курс интенсивной переподготовки и повышения квалификации. Через месяц с данной группой работников было проведено повторное тестирование и оказалось, что у 3,5% данной группы уровень квалификации все-таки остался на прежнем уровне либо вырос несущественно. Сравните доли работников, не прошедших аттестацию в прошлом месяце и в текущем. Исходя из этого, наметьте кадровые и производственные перспективы предприятия.

6. Дайте определения: перекрестная классификация, перекрестная таблица (матрица), одномерное распределение, двумерное, многомерное распределение, независимые переменные, временной ряд, кумулята и кумулятивный временной ряд.

7. Что такое классификация и группировка? Объясните место и роль метода классификации и группировки в социологическом исследовании.

8. Какие задачи в исследовании совокупностей не могут быть решены с помощью простой группировки? Каковы разновидности сложной группировки?

9. В каких случаях используются неравные интервалы? Какой вид группировки при этом предпочтителен?

10. Какие типы статистических таблиц вам известны? Почему статистическая таблица должна быть легко обозримой и иметь небольшие размеры?

11. Дайте определения: гистограмма, полигон распределения, правило нулевого начала, разрыв шкалы, многозначный

график, арифметическая шкала, арифметическая временная диаграмма, график отношений, кумулята, диаграмма полос; круговая диаграмма.

12. В чем заключается необходимость построения графических изображений при обработке социологической информации?

13. Выберите наиболее подходящий тип графика и графически представьте данные следующих таблиц:

а)

Пол студентов	Мужской	Женский
Доля в общем количестве	42%	58%

б)

Учебный год	Количество студентов
1998–1999	600
1999–2000	640
2000–2001	690
2001–2002	760
2002–2003	850

с)

Учебный год	Количество студентов		
	Муж.	Жен.	Всего
1998–1999	270	330	600
1999–2000	300	340	640
2000–2001	340	350	690
2001–2002	380	370	760
2002–2003	440	410	850

Какие выводы можно сделать на основании анализа графиков **б** и **с**? Подумайте и прокомментируйте, как в целом можно проанализировать графики, какие важные выводы можно сделать на их основе?

14. Дайте определения: характеристики положения, среднее арифметическое (простое и взвешенное), квантили,

медиана, медианный интервал, квартиль, дециль, центиль, мода, модальный интервал, модальная частота.

15. Каково будет комбинированное среднее двух групп, если среднее из 100 событий первой будет равно 10, а среднее из 50 событий второй будет равно 15? Каково было бы комбинированное среднее, если бы каждая группа состояла из 50 событий? Из 100 событий?

16. Для приведенного ниже ряда рассчитайте среднее арифметическое (простое и взвешенное), моду и медиану. Есть ли различия в данных показателях? Если есть, то как можно их объяснить?

8 8 7 5 2 8 7 8 9 8 5 8 2 6 8 2 8 4 1 8

17. В течение двух недель в кинотеатре посчитывали количество зрителей и получили следующие данные, представленные в табл. 1

Таблица 1

Данные посещаемости кинотеатра

День недели	Утренние и дневные сеансы	Вечерние сеансы
<i>1-я неделя</i>		
<i>Понедельник</i>	376	520
<i>Вторник</i>	420	560
<i>Среда</i>	397	590
<i>Четверг</i>	440	558
<i>Пятница</i>	432	614
<i>Суббота</i>	668	748
<i>Воскресенье</i>	684	711
<i>2-я неделя</i>		
<i>Понедельник</i>	567	712
<i>Вторник</i>	611	769
<i>Среда</i>	581	787
<i>Четверг</i>	645	765
<i>Пятница</i>	641	837
<i>Суббота</i>	892	961
<i>Воскресенье</i>	878	930

По приведенным в таблице данным рассчитайте следующие показатели: среднее количество зрителей в день и в неделю, в будние и выходные дни, на утренних, дневных и вечерних сеансах, в первую и во вторую неделю. Проанализируйте полученные результаты и сделайте выводы: как изменяется количество зрителей в зависимости от дней недели и времени сеансов, отличается ли количество зрителей в первую и вторую неделю и т. п.

18. По данным Таблицы 1 для каждой недели рассчитайте следующие показатели вариации: дисперсию, среднее квадратическое отклонение, коэффициент вариации. На этой основе сделайте выводы, которые могут заинтересовать администрацию кинотеатра, оказаться полезными в управлении данным заведением.

19. Таблица 2 показывает процент произведений каждого из шести композиторов в репертуаре (например, Киевского оркестра, Харьковского оркестра и т. д.). Вычислите медианный процент для каждого композитора, среднее линейное отклонение, основанное на медиане, и V для каждого композитора. Классифицируйте композиторов в соответствии с их популярностью и интерпретируйте результат.

Таблица 2

**Показатели популярности выдающихся композиторов
в разных городах Украины**

Композитор	Оркестры городов							
	Киев	Харьков	Львов	Одесса	Донецк	Ивано-Франковск	Днепропетровск	Симферополь
Бах	0,8	4,9	2,0	1,8	2,3	2,9	3,3	1,4
Бетховен	9,6	10,9	8,1	11,0	11,4	10,3	10,7	9,3
Берлиоз	2,5	1,6	1,1	0,7	2,6	2,3	0,6	0,1
Брамс	9,2	9,4	7,4	11,2	11,4	9,8	10,8	11,9
Прокофьев	1,2	1,2	0,5	1,1	1,1	1,1	1,8	1,0
Чайковский	4,5	5,9	7,3	9,6	2,8	5,4	6,5	9,3

20. Объясните, почему среднее квадратическое отклонение, а не вариация, обычно используется как мера дисперсии.

21. Если средний возраст студентов вуза равен 20 годам, а среднее квадратическое отклонение равно 2, то каким будет среднее и сигма (σ) этой группы двадцать лет спустя? Чему будет равен V ?

22. Население города состоит из 50% мужчин и 50% женщин; 70% – украинцев, 30% россиян. Можно ли представить переменные одним V_{σ} ? Обоснуйте ответ.

23. Дайте определения: коэффициент качественной вариации, подклассификация; стандартизация, скрытый фактор, причинный фактор, максимум различий, операция нормировки, отношение, процентное отношение, степень, проблемная переменная, переменная нормировки, числовая основа, открытая группа, индекс, ожидаемая величина.

24. Является ли V_q стандартизированной мерой? Аргументируйте ответ.

25. В XVIII веке европейские церкви вели записи смертей и рождений, которые впоследствии были использованы учеными для оценки размеров городов. Так, в Берлине приблизительно в 1700 году произошло около 178 смертей. Оцениваемое отношение числа смертей к общему населению равно 1:35. Вычислите степень смертности на 1000. Оцените размер города Берлина в то время.

Раздел III. ПОНЯТИЯ И ПРИНЦИПЫ СТАТИСТИЧЕСКОГО ВЫВОДА КАК ОСНОВА АНАЛИЗА СОЦИОЛОГИЧЕСКИХ ДАННЫХ

3.1. Кривые распределения, законы распределения

Понятие нормального частотного распределения. Можно сказать, что основной целью анализа вариационных рядов является выявление закономерности распределения (исключая при этом влияние случайных для данного распределения факторов). Если попробовать каждое значение признака отобразить в виде точки на плоскости координат и соединить все эти значения плавной линией, мы получим некую кривую, которая для полигона частот будет являться некоторым пределом. Эту линию называют кривой распределения [1, с. 115–119; 10].

Кривая распределения есть графическое изображение в виде непрерывной линии изменения частот в вариационном ряду, которое функционально связано с изменением вариантов. Кривая распределения отражает закономерность изменения частот при отсутствии случайных факторов. Графическое изображение облегчает анализ рядов распределения [1, с. 138–144].

Любое эмпирические распределения сравнивают с неким эталоном – идеальным распределением. Это сравнение необходимо для:

1) возможности спрогнозировать дальнейшее поведение и развитие того или иного феномена, в случае если различия между эмпирическими и теоретическими распределениями невелики;

2) выявления причин, влияющих на проявление отличий между теоретическими и эмпирическими распределениями, если таковые наблюдаются.

Известно достаточно много форм кривых распределения, по которым может выравняться вариационный ряд, но в практике

статистических исследований наиболее часто используются такие формы, как *нормальное распределение* и *распределение Пуассона*.

Нормальное распределение зависит от двух параметров: среднего арифметического $M[X]$ и среднего квадратического отклонения σ . Можно даже сказать, что одно из наиболее важных применений последнего и состоит в описании нормального распределения, то есть «нормальности» распределения, если можно так выразиться.

Применение количественных методов, вошедших в практику социологических исследований, в той или иной степени опирается на предположение, что изучаемый признак (или совокупность признаков) подчиняется определенному статистическому закону распределения. Таким наиболее часто встречающимся распределением и является нормальный (гауссовский) закон распределения. Нормальное распределение признака наблюдается в тех случаях, когда на величину его значений влияет множество случайных независимых или слабо зависимых факторов, каждый из которых играет в общей сумме примерно одинаковую и малую роль (то есть отсутствуют доминирующие факторы).

Функция плотности нормального распределения имеет вид

$$f(x) = \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{1}{2} \left(\frac{x - M[X]}{\sigma} \right)^2},$$

где • $M[X]$ или – среднее значение признака (математическое ожидание);

• σ – среднее квадратическое отклонение.

Несколько фактов из истории. Закон нормального распределения впервые был выведен в 1783 г. Абрахамом де Муавром, английским математиком французского происхождения. Изначально наиболее активно нормальное распределение использовалось в астрономии, для описания распределения ошибок измерений вокруг «истинных» значений (точных координат) расположения того или иного небесного тела. Поэтому по сегодняшний день его часто называют кривой ошибок. В отношении

к социальным, психологическим и антропометрическим данным закон нормального распределения был впервые применен Адольфом Кетле, бельгийским статистиком, математиком, астрономом, метеорологом, социологом, в 1830 году.

По мнению А. Кетле¹⁵, социальные и моральные данные стремятся расположиться на нормальной кривой, которая является проявлением всеобщего закона природы. Изображение этой кривой приведено на графике, представленном чуть ниже (см. Рис. 3.1).

Смысл графика заключается в том, что все характеристики живой и неживой материи в этом мире распределяются по этому закону. Если, например, глубоко задуматься о том, как все люди в мире распределены по росту, взглянув на график, сразу же можно это увидеть. Становится понятным, что в мире существует три-пять процентов самых высоких людей и три-пять процентов самых низкорослых, рост всех остальных людей равномерно распределяется между этими двумя противоположностями, формируя понятие среднего.

То же самое касается и всех остальных количественных и качественных характеристик «человеческой» совокупности. Например, способности человека, те или иные особенности его поведения, склонности, черты характера, «в идеале» должны быть подчинены закону нормального распределения. Причем считается, что этот закон начинает действовать в группе людей, численность которой составляет более двух человек. Однако чем больше людей, тем нагляднее закон себя проявляет.

Вполне правдоподобным представляется тот факт, например, что люди по уровню одаренности талантом распределены согласно нормальному закону: в мире существует три-пять процентов талантливых людей и три-пять процентов «бездарей», которые, так сказать, друг друга «уравновешивают», а все остальные люди располагаются где-то между ними. Однако, если речь идет непосредственно о социологических данных, то было бы оши-

¹⁵ **Адольф Кетле** (*Ламбер Адольф Жак Кетеле*; 22 февраля 1796, Гент – 17 февраля 1874, Брюссель) – бельгийский математик, астроном, метеоролог, социолог. Один из родоначальников научной статистики.

бочным считать распределения, отличающиеся от нормального, необычными или «незаконными». Взять, к примеру, картину распределения населения Украины по уровню материального положения: было бы крайне оптимистичным полагать, что в нашем государстве всего лишь 3–5% тех, кто находится у черты бедности. По данным Института демографии и социальных исследований им. Птухи НАН Украины, после экономического кризиса 2008–2009 годов в Украине стало больше людей, ощущающих себя нищими. Согласно социологическим исследованиям, проведенным Центром Разумкова¹⁶ в последние годы, уровень нищеты (то есть людей, которые едва сводят концы с концами и их дохода не хватает даже на питание) увеличился с 12,5% в начале 2008 года до 13,8% в октябре 2011 года. Количество бедных людей, то есть тех, доходов которых хватает исключительно на питание и вещи самой первой необходимости, соответственно также увеличилось с 34% до 41%. Это данные – основанные на самооценках граждан Украины. По объективным показателям дохода и имущественного признака мы имеем 10% самых бедных украинцев и 2–5% самых богатых. Причем разница в уровне их материальной обеспеченности составляет (на момент конца 2001 – начала 2012 гг.) от 40 до 60 раз (с учетом того, что «нормальный» разрыв между «бедняками» и «богачами» должен быть не более, чем в 4–6 раз) [17].

Распределение необработанных эмпирических данных, полученных в результате социологического исследования, может подчиняться любому случайному распределению. Тем не менее как статистическая модель вариаций нормальная кривая наиболее популярна, несмотря на то, что она не обладает универсальностью, приписанной ей А. Кетле.

Характеристики нормальной кривой. По определению, нормальная кривая *состоит из бесконечного числа точек, унимодальна, симметрична и не ограничена* в обоих направ-

¹⁶ Украинский негосударственный аналитический центр, который проводит исследования государственной политики в различных сферах деятельности (внутренняя и внешняя политика, экономическая политика, социальная политика, государственное управление и др.)

лениях. Следовательно, *мода, медиана и среднее арифметическое равны* по величине и делят распределение на две равные части. Графическое изображение этого распределения представляет собой ровную, колоколообразную кривую с характерным крылом, которое нигде не касается базовой линии. Начиная с максимума, кривая спадает все сильнее до точек перегиба, а затем становится более пологой, простираясь бесконечно в обоих направлениях. Точки перегиба находятся на расстоянии, равном одной сигме от максимума, делящего кривую на две симметричные части. Графически сигма равна линейному расстоянию вдоль базовой линии от центра до координаты, определяющей точки перегиба.

Каждому значению сигмы – расстоянию от среднего – строго соответствует определенный процент площади или частоты; следовательно, «сигма-точки» служат удобной мерой положения. Интервал переменной в одну сигму, отсчитываемый от среднего, соответствует 34%. Следующий интервал между средним и точкой, соответствующей двум сигмам (2σ), имеет относительную частоту не 68%, а лишь 48 из-за изменения наклона кривой при увеличении расстояния от среднего. Так как распределение симметрично, то приблизительно две трети случаев (68,26%) заключены внутри центрального интервала, который простирается от $-\sigma$ до $+\sigma$. 95% случаев попадают внутрь интервала размером в две сигмы по обе стороны от среднего и практически 100% (99,72%) попадают внутрь интервала размером в три сигмы по обе стороны от среднего.

Аналогично можно определить долю случаев между средним и любым другим значением интервала на основной линии, выраженном в единицах сигмы. Для этой цели составлены справочные таблицы. Наиболее известная из них – таблица нормальных площадей, заключенных под кривой между средним и выбранным значением переменной в единицах сигм – которая содержит долю полной площади. Такая таблица в сокращенном виде приводится в конце учебника (см. Приложение 2, *табл. Е*). Она описывает лишь одну половину распределения, так как последнее симметрично.

Изучение таблицы показывает, что приблизительно 0,4953 случаев расположены в пределах интервала, длиной $2,6\sigma$, отсчитываемого от среднего, то есть практически 99% лежат внутри интервала в $2,6\sigma$ по обе стороны от среднего. Следовательно, для всех практических целей в случае нормального распределения достаточно иметь интервал переменной длиной в 6σ .

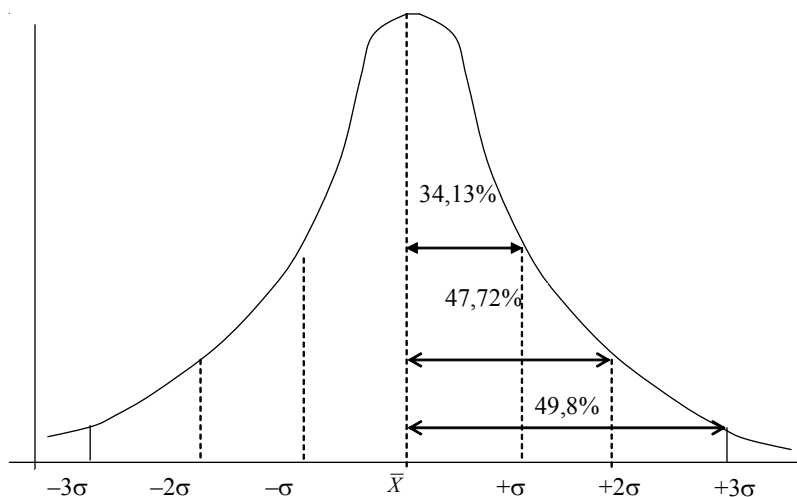


Рис. 3.1. Кривая нормального распределения

Другие законы «идеальных» распределений. К другим типам распределений, реже встречающимся и поэтому не рассматриваемым подробно в данном издании, можно отнести:

1. Закон распределения Пуассона. Определить его можно в случае, если значение среднего арифметического близко к значению среднего квадратического отклонения. По закону Пуассона развивается процесс спроса и предложения.

2. Биномиальный закон распределения, который применяется для альтернативных, качественных признаков.

Стандартное и нормированное отклонение. Выбор сигмы в качестве единицы измерения на базовой линии нормального распределения дает единицу, независимую не только от различных систем измерения, но также и от самих конкретных значений.

Для нее безразлично, имеем ли дело с доходами в сотню или миллион долларов, с длительностью в десять секунд или десять лет, или с изменяющейся напряженностью установки.

Однако первичные данные всегда поступают к исследователю в необработанном виде, и поэтому нельзя, например, быстро сравнить оклады учителей и их стаж работы, даже если обе переменные имеют нормальные распределения. Решение этой проблемы лежит, конечно, в переводе первичных данных в сигма-единицы, то есть в сравнимые показатели.

С этой целью отклонения от соответствующего среднего представляются как частные от деления интервалов на сигму. Полученные показатели называются нормированными отклонениями. Их принято обозначать через Z .

$$Z = \frac{x_i - M[X]}{\sigma},$$

где • x_i значение признака;

• $M[X]$ – среднее значение;

• σ – среднее квадратическое отклонение.

Нанеся абсолютные Z и σ шкалы на базовую линию, можно наглядно показать эквивалентность первичных данных нормированными отклонениями. Более того, с помощью наложения шкал можно обнаружить скрытое первоначально тождество между нормально распределенными переменными.

Для превращения любых последовательностей значения в стандартные отклонения, прежде всего, необходимо вычислить среднее и σ . Перевод в сигма-единицы позволяет установить относительное положение каждого события в совокупности путем использования таблицы нормальных площадей или частот. Из этой таблицы, например, определяется, что точка $1,5\sigma$, отсчитанная от среднего, соответствует приблизительно 43% полной частоты и т. д. Необходимо подчеркнуть, что, хотя величина сигм может быть вычислена для любого распределения, вышеприведенный перевод в сигма-единицы справедлив лишь для нормального распределения.

3.2. Проверка нормальности распределения

Так уж устроены мы, люди, что нами «все познается в сравнении». Не зная, что такое «плохо», мы не в состоянии определить, что такое «хорошо». Наиболее часто в качестве эталона для сравнения мы выбираем некий идеал, как правило, не нами придуманный, существующий объективно. Для того чтобы осуществить анализ данных, полученных в ходе социологического исследования, для того чтобы эти данные превратить в достоверную информацию, социолог также должен прибегнуть к процедуре сравнения. За эталон сравнения в данном случае принимается «идеальное» нормальное распределение. Самая важная социологическая информация получается именно в результате подобного сравнения, на основе определения того, насколько сильно и по каким параметрам эмпирическое распределение значений того или иного признака отличается от этого идеала.

В арсенале математической статистики существует множество методов сравнения эмпирических распределений с «идеальной» теоретической моделью нормального распределения. Вообще, принадлежность наблюдаемых данных нормальному закону является необходимой предпосылкой для корректного применения большинства классических методов математической статистики, используемых в процедуре измерения. Поэтому проверка на отклонение от нормального закона является частой процедурой и при анализе социологических данных [14].

Начинается эта процедура с выдвижения статистической гипотезы, продолжается выбором критерия проверки этой гипотезы и завершается вычислением коэффициента, который свидетельствует о наличии либо отсутствии, силе и/или направлении связи.

Статистический вывод. Проверка гипотез. Так как во многих практических задачах точный закон распределения неизвестен, он представляет собой гипотезу, которая требует статистической проверки. Статистические гипотезы могут иметь различные формулировки, но по своей сути они являются гипотезами о распределении той или иной случайной величины.

На практике часто приходится делать некоторые выводы по имеющемуся у нас небольшому объему данных (выборки) о свойствах всей генеральной совокупности. Эти выводы осуществляются с помощью определенных статистик и поэтому называются статистическими. Теория статистического вывода занимает центральное место в статистике. Основным способом, с помощью которого делаются статистические выводы, является проверка гипотез [3, с. 31–33; 5; 9].

Существует два вида гипотез: 1) научные и 2) статистические. *Научная гипотеза* – это предполагаемое решение некоторой проблемы. Она обычно формулируется в виде теоремы. **Статистическая гипотеза** – некоторое утверждение относительно неизвестного параметра или какой-либо характеристики. Например, среднее значение генеральной совокупности равно 125 ($M[X]=125$) или коэффициент корреляции равен 0 ($r=0$).

Для проверки статистических гипотез используются статистические критерии, которые представляют собой некоторое правило, по которому мы делаем вывод о правильности или неправильности рассматриваемой статистической гипотезы [3, с. 31–32; 10]. Подробнее о статистических критериях будет сказано в подразделах 3.3 и 3.4.

При использовании этих критериев возможны четыре исхода проверки гипотез:

1. Гипотеза H_0 верна и принимается, согласно критерию.
2. Гипотеза H_0 неверна и отвергается, согласно критерию.
3. Гипотеза H_0 верна, но отвергается, согласно критерию (*ошибка первого рода*).
4. Гипотеза H_0 не верна, но принимается, согласно критерию (*ошибка второго рода*).

Как видим, при проверке статистических гипотез всегда существует вероятность допустить ошибку, приняв неверную или опровергнув верную гипотезу. При этом *разграничивают ошибки первого или второго рода*.

Ошибкой первого рода называется ошибка, состоящая в опровержении верной гипотезы, когда делается неверное заключение о том, что гипотеза не соответствует истине, когда в

действительности такое соответствие есть. При этом, согласно критерию, отвергается гипотеза H_0 , которая на самом деле верна. Вероятность такой ошибки не выше определенного уровня значимости. **Уровнем значимости α** называется вероятность совершения ошибки первого рода. Значение уровня значимости обычно задается близким к нулю (например, 0,05; 0,01; 0,02 и т. д.), потому что, чем меньше это значение, тем меньше вероятность совершения ошибки первого рода, состоящей в опровержении верной гипотезы H_0 .

Ошибкой второго рода называется ошибка, состоящая в принятии ложной гипотезы, когда делается неверное заключение, о том, что гипотеза соответствует истине, когда в действительности такого соответствия нет. Вероятность совершения ошибки второго рода, т. е. принятия ложной гипотезы, обозначается как β .

При проверке нулевой гипотезы возможно возникновение следующих ситуаций:

H_0	Верная	Ложная
Отклоняется	Ошибка первого рода	Решение верное
Не отклоняется	Решение верное	Ошибка второго рода

По сути, проверка статистической гипотезы означает проверку согласования исходных выборочных данных с выдвинутой основной гипотезой. **Общая схема проверки статистической гипотезы** состоит из пяти этапов:

I – *выдвигаем две статистические гипотезы:*

- 1) основную – нулевую (H_0),
 - 2) конкурирующую – альтернативную (H_1).
- Например, H_0 среднее значение ГС = 125.
 H_1 среднее значение ГС $\neq$ 125.

Математически это можно записать:

$$H_0: M[X] = 125,$$

$$H_1: M[X] \neq 125 \quad (M[X] < 125; M[X] > 125).$$

II – *задаемся уровнем значимости.* Статистический вывод никогда не может быть сделан со стопроцентной уверенностью.

Всегда допускается риск принятия неправильного решения. При проверке статистических гипотез мерой такого риска и выступает уровень значимости, который обычно обозначается. Фактически уровень значимости представляет собой долю и процент ошибок, которые мы можем себе позволить при статистических выводах. Чаще всего используют следующие три значения уровня значимости: $\alpha = 0,1$ или 10%; $\alpha = 0,05$ или 5%; $\alpha = 0,01$ или 1%. Наиболее популярным из них является $\alpha = 0,05$ или 5% (допускается 5% ошибка на 100% выборки).

III – по исходным данным, то есть по выборке, *вычисляем наблюдаемое значение статистического критерия*. В общем случае обозначим его как $K_{набл}$.

IV – по специальным статистическим таблицам (которые в данном издании приводятся в приложениях) *находим критическое значение критерия* ($K_{кр}$), в зависимости от заданного уровня значимости (α) и числа степеней свободы (df).

V – *путем сравнения* найденных наблюдаемых значений критерия ($K_{набл}$) с теоретическим (табличным) критическим значением ($K_{кр}$) делаем вывод о правильности/неправильности, принятии/отвержении гипотезы.

- Если найденное для выборки $K_{набл}$ меньше табличного $K_{кр}$, то гипотеза H_0 принимается на заданном уровне значимости α . В этом случае наблюдаемое различие двух сравниваемых нами совокупностей объясняется только случайностью выборки. Хотя принятие гипотезы H_0 совсем не означает, что параметры этих совокупностей полностью одинаковы. Считается, что имеющийся в распоряжении статистический материал просто не дает оснований для отклонения данной гипотезы.

- Если же найденное для выборки $K_{набл}$ больше табличного $K_{кр}$, то гипотеза H_0 отклоняется в пользу гипотезы H_1 при заданном уровне значимости α . В этом случае наблюдаемое отличие сравниваемых совокупностей нельзя объяснить только лишь случайностью, скорее всего, имеет место наличие одного или нескольких весомых факторов, влияющих на распределение.

В определенном смысле задача проверки статистических гипотез логически предшествует задаче оценивания. Например,

при статистическом исследовании имеющейся разницы между средними двух нормальных совокупностей исследователю приходится установить, является ли данная разница статистически значимой, не является ли она результатом выборочной случайной флуктуации, то есть проверить гипотезу о равенстве средних. Если эта гипотеза не подтвердится, исследователю, возможно, придется перейти к следующему этапу – оцениванию величины разницы между средними.

Необходимые для проверки различных гипотез вычисления могут быть сделаны с помощью специально разработанных компьютерных статистических программ. Однако важно понимать, что истинным критерием проверки статистических гипотез все-таки остается опыт самого исследователя. По этому поводу согласимся с учеными, утверждающими, что статистические критерии значимости – лишь формально точный инструмент. Чем больше социолог знает об исследуемом им явлении, тем точнее будет сформулирована гипотеза и выводы, сделанные на основе критериев значимости [11; 12; 15].

3.3. Общие представления о критериях «нормальности» распределения

Как уже говорилось, тестирование данных на нормальность является необходимым для их полноценного, полного анализа, так как большое количество статистических методов исходит именно из предположения нормальности распределения изучаемых данных. Существует группа статистических критериев, предназначенных для проверки нормальности распределения.

Статистический критерий (K) – специально подобранная случайная величина, точное либо приближенное распределение которой известно и обычно сведено в таблицы.

Множество всех возможных значений выбранного статистического критерия делится на два непересекающихся подмножества. Первое подмножество включает в себя те значения критерия, при которых основная гипотеза отвергается, а второе

подмножество – те значения критерия, при которых основная гипотеза принимается.

Критической областью называется множество возможных значений статистического критерия, при которых основная гипотеза отвергается. Областью принятия гипотезы или областью допустимых значений называется множество возможных значений статистического критерия, при которых основная гипотеза принимается. Если наблюдаемое значение статистического критерия, рассчитанное по данным выборочной совокупности, принадлежит критической области, то основная гипотеза отвергается. Если наблюдаемое значение статистического критерия принадлежит области принятия гипотезы, то основная гипотеза принимается.

Критическими точками или квантилями ($K_{кр}$) называются точки, разграничивающие критическую область и область принятия гипотезы. Критические области могут быть как односторонними (право- или левосторонними), так и двусторонними.

Правосторонняя критическая область характеризуется неравенством вида:

$$K_{набл} > K_{кр},$$

где • $K_{набл}$ – наблюдаемое значение статистического критерия, вычисленное по данным выборки;

• $K_{кр}$ – положительное критическое (теоретическое) значение статистического критерия, определяемое по таблице распределения данного критерия. Следовательно, для определения правосторонней критической области необходимо рассчитать положительное значение статистического критерия $K_{кр}$.

Левосторонняя критическая область характеризуется неравенством вида:

$$K_{набл} < K_{кр},$$

где • $K_{набл}$ – эмпирическое (наблюдаемое) значение статистического критерия, вычисленное по данным выборки;

• $K_{кр}$ – отрицательное критическое (теоретическое) значение статистического критерия, определяемое по таблице распределения данного критерия.

Следовательно, для определения правосторонней критической области необходимо рассчитать отрицательное значение статистического критерия $K_{кр}$.

Двусторонняя критическая область характеризуется двумя неравенствами вида:

$$K_{набл} > K_{кр1} \text{ и } K_{набл} < K_{кр2}, \text{ при этом } K_{кр1} > K_{кр2},$$

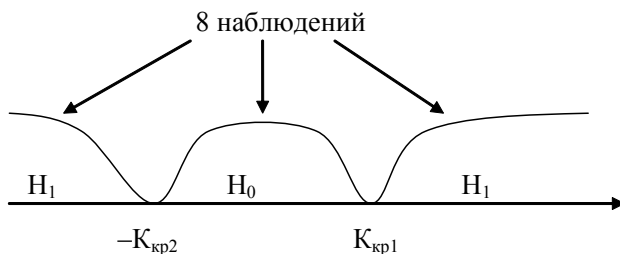
где • $K_{набл}$ – наблюдаемое значение статистического критерия, вычисленное по данным выборки;

• $K_{кр1}$ – положительное значение статистического критерия, определяемое по таблице распределения данного критерия;

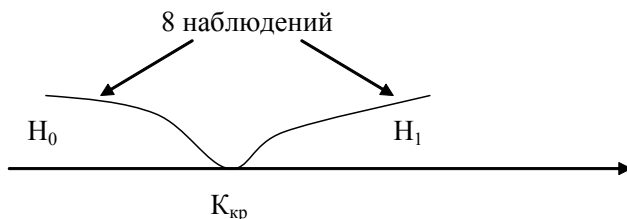
• $K_{кр2}$ – отрицательное значение статистического критерия, определяемое по таблице распределения данного критерия.

Проиллюстрируем наглядно наиболее часто встречававшиеся примеры:

А) Для двусторонней критической области



Б) Для односторонней критической области



С учетом того, что в случае принятия нулевой гипотезы дисперсии сравниваемых распределений будут равны, то есть

$H_0: D(X) = D(Y)$, выбор критической области осуществляется исходя из вида конкурирующей гипотезы $H1$. При этом:

1) *правосторонняя критическая область* выбирается в том случае, если $H1: D(X) > D(Y)$;

2) *левосторонняя критическая область* выбирается в том случае, если $H1: D(X) < D(Y)$;

3) *двусторонняя критическая область* выбирается в том случае, если $H1: D(X) \neq D(Y)$.

Основной вопрос, возникающий на основе сказанного, заключается в определении того, ***какой из критериев, односторонний или двусторонний, выбрать в том или ином случае.*** Ответ на этот вопрос следует искать за пределами математико-статистического формализма, он в полной мере зависит от целей и задач социологического исследования. Если логикой исследования допускается, что различия сравниваемых параметров могут принимать как положительные, так и отрицательные значения, то следует применять двусторонний критерий. Если же есть дополнительная информация (основанная, например, на опыте предыдущих исследований) о том, что один параметр может принимать значения либо больше, либо меньше значений другого параметра, то используется односторонний критерий. Причем, если имеются основания, достаточные для применения одностороннего критерия, то следует именно его (а не двусторонний) и предпочесть, так как односторонний критерий полнее использует информацию об изучаемом явлении либо процессе и поэтому чаще дает правильные и более точные результаты [12].

Мощностью статистического критерия называется вероятность попадания данного критерия в критическую область, при условии, что справедлива конкурирующая гипотеза $H1$, то есть выражение $1-\beta$ является мощностью критерия. Ясно, что мощность может принимать любые значения от 0 до 1. Чем ближе мощность к единице, тем эффективнее критерий.

Критическую область следует строить так, чтобы мощность критерия была максимальной. Повторимся, что для выполнения $\alpha = 0,05$ или 5% (т.е. допускается 5% ошибки на 100% выборки). Но если выводы, которые предстоит сделать по результатам

проверки гипотез, связаны с большой ответственностью, то следует выбирать $\alpha = 0,01$, или $\alpha = 0,001$. Это обеспечивает минимальную ошибку второго рода, состоящую в том, что будет принята неправильная гипотеза. Однако вполне резонным является вопрос о том, *как трактовать результаты исследования, связанные с установлением того или иного уровня значимости*. Часто поступают следующим образом: уровень значимости не устанавливается точно, а по экспериментальным данным вычисляется вероятность (P) того, что критерий $K_{набл}$ выйдет за пределы $K_{кр}$. Окончательные результаты обычно приводят в следующем виде: 1) если вычисленное значение критерия не превосходит критического (табличного) значения на уровне $\alpha = 0,05$, то различие считается статистически незначимым; 2) если вычисленное по выборке значение превышает критические значения $\alpha = 0,05$, $\alpha = 0,01$ и $\alpha = 0,001$, то результаты записываются следующим образом: соответственно, $p < 0,05$; $p < 0,01$; $p < 0,001$. Это значит, что наблюдаемые различия исследуемых совокупностей статистически значимы на уровнях значимости 0,05 или 0,01, или 0,001.

Критерии значимости подразделяются на три типа:

1. Параметрические критерии значимости, то есть те, которые служат для проверки гипотез о параметрах распределения генеральной совокупности (чаще всего – нормального распределения). *Параметрические критерии включают в формулу расчета параметры распределения, то есть средние и дисперсии.*

2. Непараметрические критерии значимости, то есть те, которые для проверки гипотез не используют предположений о параметрах распределения (например, в случаях, когда эти параметры неизвестны). *Непараметрические критерии не включают в формулу расчета параметров распределения и основаны на оперировании частотами или рангами.*

3. Критерии согласия – как отдельная, особая группа, которую составляют критерии, служащие для проверки гипотез о согласованности генеральной совокупности, из которой получена выборка, с ранее принятой теоретической моделью (чаще всего – моделью нормального распределения). Именно поэтому

критерии согласия часто ассоциируют с так называемыми критериями «нормальности».

Вообще же, критерии нормальности являются лишь частным случаем критериев согласия. Однако именно к этим критериям наиболее часто прибегают социологи при осуществлении анализа эмпирических данных.

В данном издании обратимся к подробному рассмотрению некоторых из этих критериев, представляющих особую важность для исследователя-социолога. Речь идет, прежде всего, о таких критериях, как: • *критерий асимметрии* (параметрический и/или согласия), • *критерий Пирсона Хи-квадрат* (непараметрический и/или согласия), • *критерий Стьюдента* (параметрический), • *критерий Фишера* (параметрический). О них и пойдет речь в следующем, заключительном подразделе этого раздела.

3.4. Критерии «нормальности»: принципы нахождения и расчетные формулы

Предварительная проверка нормальности распределения: критерий асимметрии (A_s). Все рассматриваемые нами в данном издании критерии значимости обеспечивают наивысшую степень достоверности статистических выводов только в тех случаях, если выборки получены из нормального распределения генеральной совокупности. При отклонении от нормального распределения эта степень существенно снижается. Следовательно, чтобы уверенно применять эти критерии, необходимо изначально осуществить проверку предположения о нормальном распределении генеральной совокупности. Для этого, в общем-то, и существуют критерии согласия. Но, применение большинства этих критериев – довольно трудоемкая вычислительная работа, занимающая немало времени. Поэтому целесообразным представляется осуществление предварительной проверки на нормальность с использованием более простых методов. Они, конечно, обладают сравнительно меньшей мощностью, помогают установить только довольно большие расхождения с нормальным распределением, но и этого может быть достаточно,

чтобы исследователь мог избавиться от необходимости осуществлять сложные вычислительные процедуры, сэкономив тем самым и драгоценное время, и энергию.

Такую упрощенную предварительную проверку можно осуществить с помощью критерия асимметрии. При этом H_0 : *случайная величина имеет распределение, отличное от нормального*. Если распределение нормально, то его коэффициент асимметрии (As) будет равен нулю. Так как нулевые значения коэффициента могут иметь место и для распределений, отличных от нормального, то этот критерий следует воспринимать как критерий установления отклонения от нормальности распределения, но не установления нормальности.

Различают несколько вариантов эмпирических распределений при сравнении их с теоретическими, а именно: симметрические и скошенные, примеры их полигонов были рассмотрены в параграфе 2.6 (см. Рис. 2.10, с. 108). *Симметрическое распределение* – это распределение, где частоты равномерно удалены от среднего арифметического ($M[X]$). Однако при анализе реально существующих социальных процессов и явлений чаще встречаются *скошенные (асимметрические) распределения*. Поэтому при характеристике эмпирических рядов распределений рассматривают • *левостороннюю*, • *правостороннюю*, • *умеренно скошенную* и • *крайне скошенную асимметрии*.

При симметрическом распределении, имеет место уменьшение частот от максимального, модального значения (Mo) одинаково в обе стороны. *Умеренно-асимметрическое* распределение характеризуется тем, что большая часть частот ответов респондентов приходится на одну сторону от максимального значения. Именно такие распределения чаще всего встречаются в социологических исследованиях. В некотором роде аномальными видами распределений являются:

– *T-образное или крайне скошенное* – это распределение, в котором наибольшая численность групп находится на одном конце амплитуды колебаний;

– *V-образное распределение* – это распределение, при котором наибольшая численность групп находится на концах амплитуды

колебаний, а его минимум приходится на центральный интервал. При появлении такого распределения чаще всего его разбивают на две части и рассматривают два независимых процесса. Разбиение производится таким образом, чтобы полученные распределения приблизительно соответствовали требованию $\bar{x} = Mo = Me$.

Для того чтобы определить величину скошенности или асимметрии, вычисляют коэффициент As по следующей формуле:

$$As = \frac{M[X] - Mo}{\sigma},$$

где • $M[X]$ – среднее арифметическое значение признака;

• Mo – модальное значение;

• σ – среднее квадратическое отклонение.

Свойства асимметрии:

• коэффициент As всегда изменяется только в пределах от -3 до $+3$. Чем ближе As к граничным значениям (-3 , $+3$), тем больше скошенность;

• чем больше разница между средним арифметическим и модой (медианой), тем больше асимметрия ряда;

• если значение As положительно, то говорят, что распределение вправо скошено, если отрицательно, то – влево скошено, если $As = 0$ – то асимметрии нет, то есть распределение симметрично и $\bar{x} = Mo = Me$;

• если $As > 0$, то кривая распределения имеет длинный правый «хвост», то есть налицо правосторонняя асимметрия. При этом выполняется соотношение $Mo < Me < M[X]$;

• если $As < 0$, то асимметрия левосторонняя, кривая распределения имеет длинный левый «хвост». При этом $M[X] > Me > Mo$;

• на практике асимметрия считается значительной, если коэффициент асимметрии превышает по модулю 0,25.

Критерий Пирсона Хи-квадрат (X^2). Наиболее часто употребляемым критерием оценки того, в соответствии с каким законом распределены значения признака, является критерий Хи-квадрат (далее по тексту – X^2). Популярность обусловлена

главным образом тем, что применение его не требует предварительного знания закона распределения изучаемого признака. Кроме того, признак может принимать как непрерывные, так и дискретные значения, причем измеренные хотя бы на номинальном уровне.

Критерий χ^2 – статистический критерий для проверки гипотезы H_0 : *наблюдаемая случайная величина подчиняется некому теоретическому закону (нормального распределения)*. Следовательно, подтверждение данной гипотезы непременно требует сравнения эмпирических частот (то есть тех, которые были получены в результате социологического исследования) и теоретических частот (то есть тех, которые могли бы быть получены, если бы данное распределение было «идеально» нормальным).

По своей сути, критерий χ^2 отвечает на вопрос о том, с одинаковой ли частотой встречаются разные значения признака в эмпирическом и теоретическом распределениях или в двух и более эмпирических распределениях.

Следовательно, критерий χ^2 Пирсона, используется:

1) для сопоставления *эмпирического* распределения признака с *теоретическим* – равномерным, нормальным или каким-то иным;

2) для сопоставления *двух, трех или более эмпирических* распределений одного и того же признака.

Принцип вычисления основывается на процедуре выравнивания вариационного ряда по кривой нормального распределения, то есть предполагает обязательное нахождение теоретических частот.

Итак, критерий χ^2 для проверки H_0 имеет вид:

$$\chi^2_{\text{набл}} = \sum_{(i)} \sum_{(j)} \frac{(n_{ij} - \tilde{n}_{ij})^2}{\tilde{n}_{ij}},$$

где n_{ij} – наблюдаемая (эмпирическая) частота, число объектов в ячейке на пересечении i -й строки таблицы сопряженности

признаков и j -го столбца, так называемая фактическая клеточная частота;

- $\tilde{n}_{ij}$ – ожидаемая по H_0 (теоретическая) частота в этой же ячейке;

- $i = 1, 2, \dots, r$ (то есть r – число строк таблицы сопряженности признаков);

- $j = 1, 2, \dots, c$ (то есть c – число столбцов таблицы сопряженности признаков).

Как видим из формулы, первое, что нам необходимо сделать, – это найти ожидаемые (теоретические) частоты.

Ожидаемые (теоретические) частоты ($\tilde{n}_{ij}$) вычисляют по формуле:

$$\tilde{n}_{ij} = \frac{r_i c_j}{N},$$

где • r_i – маргинал i -й строки;

- c_j – маргинал j -го столбца;

- N – объем выборочной совокупности.

Дальнейшая последовательность действий довольно проста, мы проиллюстрировали ее на Примере 3.4.1, представленном чуть ниже.

Заметим, что для таблицы 2×2 число степеней свободы всегда равно 1.

При расчете критерия χ^2 -квадрат Пирсона должно соблюдаться следующее условие: достаточно большим должно быть число наблюдений ($n \geq 30$). Если полученное значение $\chi^2_{набл}$ меньше критических значений, представленных в таблице, то делается вывод о том, что эти расхождения между эмпирическими и теоретическими частотами распределения могут быть случайными, формулируется предположение о близости эмпирического распределения к нормальному и принимается гипотеза H_0 . Фактически же это означает, что искомая связь между признаками если и есть, то, вероятно, она слишком слаба.

Итак, рассмотрим реальный пример применения критерия χ^2 Пирсона.

Пример 3.4.1. Предположим, мы провели социологическое исследование для известного «Модного дома». В данном случае изучение взаимосвязи между возрастом респондентов и их отношением к «классической» моде могло бы быть очень информативным и полезным для заказчика исследования. Таблица сопряженности признаков «Возраст» и «Отношение к классической моде» с абсолютными частотами двумерного распределения представлена ниже (см. *Табл. 3.1*).

Таблица 3.1

**Таблица сопряженности по признакам «Возраст»
и «Отношение к классической моде» (абсолютные частоты)**

<i>Отношение \ Возраст</i>	<i>Молодежь</i>	<i>Люди среднего возраста</i>	<i>Пожилые люди</i>	ИТОГО (маргинальные частоты по строкам – r_i для числа строк – r)
<i>Очень нравится классический стиль одежды</i>	26	13	5	44
<i>Скорее нравится, чем нет</i>	20	11	8	39
<i>Скорее не нравится, чем нравится</i>	9	10	20	39
<i>Совершенно не нравится «классика» в одежде</i>	7	10	15	32
ИТОГО (маргинальные частоты по столбцам – c_j для числа столбцов – c)	62	44	48	$n=154$

Попытаемся подтвердить/опровергнуть наличие/отсутствие взаимосвязи между признаками при помощи критерия χ^2 Пирсона. Если полученное наблюдаемое значение критерия $\chi^2_{набл}$ будет меньше критического значения $\chi^2_{кр}$, по заданному уровню значимости (α), в соответствии с числом степеней свободы (df), то гипотеза H_0 принимается и делается вывод о вероятном отсутствии связи между признаками. По сути это значило бы, что и молодые, и пожилые, и люди среднего возраста (с большой вероятностью) сделают одинаковые выборы, определяя свое отношение к классической моде. А если отличий нет, значит,

признак возраста никак не влияет на это отношение. И наоборот, если $\chi^2_{набл}$ будет больше $\chi^2_{кр}$, гипотеза H_0 отвергается в пользу альтернативной H_1 и делается вывод о вероятном наличии связи.

Итак, найдем и процедуру поиска опишем при помощи расчетной таблицы, представленной ниже.

Таблица 3.2

Расчетная таблица нахождения критерия $\chi^2_{набл}$ для двумерного распределения по признакам «Возраст» и «Отношение к моде»

N_2 n/n ячейки	n_{ij} Набл., эмп. частота	$\tilde{n}_{ij}$ Ожид. частота	$(n_{ij} - \tilde{n}_{ij})$	$(n_{ij} - \tilde{n}_{ij})^2$	$\frac{(n_{ij} - \tilde{n}_{ij})^2}{\tilde{n}_{ij}}$
1	26	17,7	8,3	68,89	3,89
2	13	12,6	0,4	0,16	0,01
3	5	13,7	-8,6	73,96	5,41
4	20	15,7	4,3	18,49	1,18
5	11	11,1	-0,1	0,01	0,00
6	8	12,2	-4,2	17,64	1,45
7	9	15,7	-6,7	44,89	2,86
8	10	11,1	-1,1	1,21	0,11
9	20	12,2	7,8	60,84	4,99
10	7	12,9	-5,9	34,81	2,7
11	10	9,1	0,9	0,81	0,09
12	15	10	5	25	2,5
					$\Sigma = 25,18(\chi^2_{набл})$

Описание вычислительных действий:

- 1) выписываем слева направо наблюдаемые частоты;
- 2) для каждой из них по соответствующей формуле вычисляем ожидаемые (теоретические) частоты, то есть те, которые могли бы стоять в ячейках таблицы в «идеальном» случае независимости признаков:

$$\langle 1 \rangle: \tilde{n}_{11} = \frac{44 \times 62}{154} = 17,7; \quad \langle 2 \rangle: \tilde{n}_{12} = \frac{44 \times 44}{154} = 12,6;$$

$$\begin{aligned}
\llcorner 3 \rrcorner: \tilde{n}_{13} &= \frac{44 \times 48}{154} = 13,7; & \llcorner 8 \rrcorner: \tilde{n}_{32} &= \frac{39 \times 44}{154} = 11,1; \\
\llcorner 4 \rrcorner: \tilde{n}_{21} &= \frac{39 \times 62}{154} = 15,7; & \llcorner 9 \rrcorner: \tilde{n}_{33} &= \frac{39 \times 48}{154} = 12,2; \\
\llcorner 5 \rrcorner: \tilde{n}_{22} &= \frac{39 \times 44}{154} = 11,1; & \llcorner 10 \rrcorner: \tilde{n}_{41} &= \frac{32 \times 62}{154} = 12,9; \\
\llcorner 6 \rrcorner: \tilde{n}_{23} &= \frac{39 \times 48}{154} = 12,2; & \llcorner 11 \rrcorner: \tilde{n}_{42} &= \frac{32 \times 44}{154} = 9,1; \\
\llcorner 7 \rrcorner: \tilde{n}_{31} &= \frac{39 \times 62}{154} = 15,7; & \llcorner 12 \rrcorner: \tilde{n}_{43} &= \frac{32 \times 48}{154} = 10.
\end{aligned}$$

Представленные выше результаты вычислений занесены в третий столбец расчетной таблицы;

3) далее находим разность между наблюдаемыми и ожидаемыми частотами;

4) полученные разности возводим в квадрат;

5) делим квадрат разностей на теоретические частоты;

6) все результаты деления складываем – эта сумма и является искомым наблюдаемым значением критерия χ^2 Пирсона.

В итоге мы получили $\chi_{набл}^2 = 25,18$. Однако само по себе это число нам ни о чем сказать не может. Полученное значение критерия мы должны сравнить с критическим значением $\chi_{кр}^2$ по таблице (см. Приложение 2, табл. Б). В таблице критическое значение χ^2 находится по заданному уровню значимости α (0,05; 0,01 или 0,001) и числу степеней свободы df , которое находится по формуле:

$$df = (r - 1)(c - 1),$$

где • r – число строк таблицы сопряженности,

• c – число столбцов таблицы сопряженности.

Число степеней свободы для нашего случая будет следующим:
 $df = (4 - 1)(3 - 1) = 6$.

По таблице критических значений χ^2 находим, что наблюдаемое значение критерия, соответствующее $df = 6$, превышает критические значения по всем уровням значимости:

$$\chi_{кр}^2 = \begin{cases} 12,59 & \text{для } p \leq 0,05 \\ 15,03 & \text{для } p \leq 0,02 \\ 16,81 & \text{для } p \leq 0,01 \end{cases}$$

$$\chi_{набл}^2 = 25,18 > \chi_{кр}^2 = 16,81.$$

Вывод: Гипотеза H_0 отвергается, распределения по разным возрастным группам различаются в своем отношении к классической моде.

Т-критерий Стьюдента (критерий, основанный на нормальном распределении). Т-критерий Стьюдента – общее название для статистических тестов, в которых статистика критерия имеет распределение Стьюдента. Распределение Стьюдента, по сути, представляет собой сумму нескольких нормально распределенных случайных величин. Чем больше величин, тем больше вероятность, что их сумма будет иметь нормальное распределение. Таким образом, количество суммируемых величин определяет важнейший параметр формы данного распределения – число степеней свободы (df).

Наиболее часто t-критерий Стьюдента используется для:

- 1) установления сходства-различия средних арифметических значений в двух независимых выборках или для установления сходства-различия двух эмпирических распределений;
- 2) установления сходства-различия двух дисперсий в двух зависимых выборках.

Ограничения в применении t-критерия Стьюдента [10; 16]:

- 1) это параметрический критерий, поэтому необходимо, чтобы распределение признака, по крайней мере, не отличалось от нормального распределения;

2) разные формулы расчета

1 для независимых выборок:

$$t_{\text{набл}} = \frac{|\bar{x}_1 - \bar{x}_2|}{\sqrt{(n_1 - 1)\delta_1^2 + (n_2 - 1)\delta_2^2}} \times \sqrt{\frac{n_1 n_2 (n_1 + n_2 - 2)}{(n_1 + n_2)}},$$

где • $\bar{x}_1$ и $\bar{x}_2$ – средние арифметические, соответственно, в 1-й и 2-й сравниваемых выборках;

- n_1, n_2 – количество респондентов в 1-й и 2-й выборках;
- δ_1, δ_2 – средние квадратические отклонения в 1-й и 2-й выборках.

В данном случае количество степеней свободы для нахождения критического значения критерия ($t_{\text{кр}}$): $df = n_1 + n_2 - 2$.

2 для зависимых выборок (с учетом корреляция результатов, поскольку измерения проводятся на одних и тех же респондентах в различных условиях (x и y)):

$$t_{\text{набл}} = \frac{(\sum d_i)}{\sqrt{\sum d_i^2 - \frac{(\sum d_i)^2}{n}}} \times \sqrt{\frac{n-1}{n}},$$

где • $d_i = x_i - y_i$, то есть разность значений признака в разных условиях (x и y), для каждого респондента;

- n – количество респондентов в выборке (объем выборки).

В данном случае количество степеней свободы для нахождения критического значения критерия ($t_{\text{кр}}$): $df = n - 1$. Проверяется статистическая гипотеза о соответствии распределения разностей t -распределению Стьюдента с нулевым средним значением.

Гипотезы:

1 для независимых выборок:

H_0 : средние значения признака в обеих выборках не различаются, линейная связь между выборками отсутствует;

H_1 : средние значения признака в обеих выборках статистически значимо различаются (гипотеза сдвига).

2 для зависимых выборок (т.е. для выборок, состоящих из одних и тех же респондентов, опрошенных в разное время):

H_0 : разности оценок респондентов в двух состояниях не отличаются от нуля;

H_1 : разности оценок испытуемых в двух состояниях статистически значимо отличаются от нуля.

И все-таки наиболее часто t -критерии применяются для проверки равенства средних значений в двух выборках. Такое сравнение может быть интересным и даже необходимым, если исследователь хочет сравнить результаты нескольких аналогичных исследований, проведенных, например, в разные годы или в разных странах. Рассмотрим на примере [9, с. 64].

Пример 3.4.2 (независимые выборки). Предположим, имеется две независимые выборки школьников, интеллект которых развивали в течение некоторого времени по двум различным методикам. Требуется установить, какая из методик лучше (см. Табл. 3.3). Предварительно было выяснено, что начальный уровень интеллекта был одинаковым в обеих выборках. Однако после применения методик средний показатель интеллекта респондентов оказался выше. Задача состоит в том, чтобы разобраться, случайны ли эти различия или закономерны, то есть зависят ли они от самой методики? Такая задача сравнения двух методик может быть переформулирована на язык статистики как задача сравнения средних арифметических значений интеллекта в обеих выборках.

Таблица 3.3

Экспериментальные данные по двум выборкам

Числовые характеристики выборки	1-я выборка	2-я выборка
n	30	32
$\bar{x}$	103	110
δ	10	12

H_0 : средние значения уровня интеллекта в обеих выборках не различаются (а следовательно, эффективность методик одинаковая);

H_1 : средние значения уровня интеллекта в обеих выборках статистически значимо различаются (а следовательно, эффективность методик разная).

Понятно, что в данном случае для получения эмпирического значения t -критерия используется формула для сравнения средних показателей двух независимых выборок, то есть формула «1».

Подставив числа, зафиксированные в таблице 3.3, в формулу $t_{\text{набл}}$ для независимых выборок, получим следующее числовое выражение:

$$t_{\text{набл}} = \frac{|103 - 110|}{\sqrt{(30-1) \times 10^2 + (32-1) \times 12^2}} \times \sqrt{\frac{30 \times 32(30+32-2)}{(30+32)}} \approx 2,486.$$

Для того чтобы сделать вывод о принятии/отвержении гипотезы H_0 , необходимо сравнить наблюдаемое (эмпирическое) значение критерия – $t_{\text{набл}}$ с критическими табличными значениями $t_{\text{кр}}$, с помощью специальной таблицы критических значений t (см. Приложение 2, табл. В). По аналогии с критерием χ^2 , для того, чтобы осуществить это сравнение, нужно найти число степеней свободы (df). В нашем примере $df = 30 + 32 - 2 = 60$.

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

Полученное наблюдаемое значение t -критерия превышает критическое для $\alpha = 0,05$, но оказывается меньше критического для $\alpha = 0,01$, то есть:

$$2,0 < 2,486 < 2,66.$$

Вывод: H_0 гипотеза отклоняется. А следовательно, можно сделать вывод о статистически значимом различии средних арифметических значений в двух выборках для $p \leq 0,05$ и о преимуществах второй методики по сравнению с первой (т.к. средний показатель интеллекта по второй выборке выше).

Пример 3.4.3 (для зависимых выборок). Допустим, проводится измерение уровня ситуативной напряженности в рабочем коллективе, состоящем из восьми человек до и после психологического тренинга, проведенного специалистом. Измерение осуществляется посредством метода анкетирования, с использо-

ванием специально разработанного, специфического инструментария. Основной вопрос, который интересует исследователя: приводит ли тренинг к изменению уровня напряженности?

Гипотезы:

H_0 : разности оценок ситуативной тревожности респондентов до и после тренинга не отличаются от нуля. Следовательно, тренинг не дает результатов;

H_1 : разности оценок ситуативной тревожности респондентов до и после тренинга статистически значимо отличаются от нуля, а следовательно, посредством данного тренинга возможно влиять на изменение уровня напряженности в коллективе.

Результаты исследования были занесены в специальную таблицу, в которую нами были добавлены дополнительные «расчётные» столбцы (см. Табл. 3.4).

Таблица 3.4

Показатели тревожности ДО и ПОСЛЕ тренинга

№ п/п респондента	Показатель тревожности до тренинга (x_i)	Показатель тревожности после тренинга (y_i)	«Расчетные» столбцы	
			$d_i = x_i - y_i$	$d_i^2 = (x_i - y_i)^2$
1	30	20	10	100
2	33	17	16	256
3	41	21	20	400
4	50	43	7	49
5	36	39	-3	9
6	45	11	34	1156
7	31	28	3	9
8	25	20	5	25
$N = 8$			$\sum d_i = 92$	$\sum d_i^2 = 2004$

Подставив в формулу «2» для нахождения t-критерия найденные значения $\sum d_i$ и $\sum d_i^2$, получим:

$$t_{\text{набл}} = \frac{92}{\sqrt{2004 - \frac{92^2}{8}}} \times \sqrt{\frac{8-1}{8}} \approx 2,798.$$

Как и в предыдущих примерах, для того чтобы осуществить

сравнение $t_{\text{набл}}$ с $t_{\text{кр}}$, необходимо найти число степеней свободы (df). В данном случае $df = 8 - 1 = 7$.

В соответствии с полученным числом степеней свободы, находим по таблице критические значения и получаем:

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

Отсюда: $2,365 < 2,798 < 3,499$.

Вывод: Принимается H_1 гипотеза. Различия в уровнях тревожности до и после тренинга следует признать статистически значимыми ($p < 0,05$), так как наблюдаемое (эмпирическое) значение превышает первое критическое (табличное), но меньше второго. Следовательно, посредством данного тренинга действительно можно добиться снижения уровня напряженности в коллективе в критических ситуациях.

Все разновидности критерия Стьюдента являются параметрическими и основаны на дополнительном предположении о нормальности выборки данных. Поэтому перед применением критерия Стьюдента рекомендуется выполнить предварительную проверку нормальности. Если гипотеза нормальности отвергается, можно проверить другие распределения, если и они не подходят, то следует воспользоваться непараметрическими статистическими тестами.

F-критерий Фишера (для сравнения дисперсий). *F-критерий Фишера* также является параметрическим критерием и используется для сравнения дисперсий двух вариационных рядов [10; 13, с. 388–389].

F-критерий Фишера используется:

- 1) в регрессионном анализе, позволяет оценивать значимость линейных регрессионных моделей. В частности, он используется для проверки целесообразности включения или исключения независимых переменных (признаков) в регрессионную модель;
- 2) в дисперсионном анализе, позволяет оценивать значимость факторов и их взаимодействия.

F-критерий Фишера основан на дополнительных предположениях о независимости и нормальности выборок данных.

Поэтому перед его применением рекомендуется выполнить предварительную проверку нормальности распределения.

Наблюдаемое (эмпирическое) значение критерия вычисляется по формуле:

$$\varphi_{\text{набл}} = \frac{\delta_1^2}{\delta_2^2},$$

где • δ_1^2 – бóльшая дисперсия;

• δ_2^2 – меньшая дисперсия рассматриваемых вариационных рядов.

С помощью данного критерия проверяется гипотеза H_0 : дисперсии сравниваемых совокупностей равны между собой, то есть

$$H_0 = \{ \delta_1^2 = \delta_2^2 \}.$$

Критическое значение критерия Фишера так же как и во всех предыдущих случаях, касающихся иных критериев, находится по специальной таблице (см. Приложение 2, *табл. Ж*), исходя из уровня значимости α и степеней свободы:

$$df_1 = n_1 - 1 \text{ (числителя);}$$

$$df_2 = n_2 - 1 \text{ (знаменателя),}$$

где n_1 и n_2 – число респондентов, составивших 1-ю и 2-ю выборки, соответственно.

Если вычисленное значение критерия $\varphi_{\text{набл}}$ больше критического $\varphi_{\text{кр}}$ для определенного уровня значимости (α) и соответствующих чисел степеней свободы для числителя и знаменателя (df_1 и df_2), то гипотеза H_0 отвергается и принимается альтернативная гипотеза H_1 ; дисперсии считаются различными.

Ниже проиллюстрировано применение критерия Фишера [13, с. 389].

Пример 3.4.4. Дисперсия такого показателя, как стрессоустойчивость для учителей составила 6,17 ($n_1 = 32$), а для менеджеров 4,41 ($n_2 = 33$). Определим, можно ли считать уровень дисперсий примерно одинаковым для данных выборок на уровне значимости $\alpha = 0,05$.

Для ответа на поставленный вопрос определим наблюдаемое значение критерия:

$$\varphi_{\text{набл}} = \frac{6,17}{4,41} \approx 1,4, \text{ при этом } df_1 = 32-1=31, \text{ а } df_1 = 33-1 = 32.$$

Исходя из проделанных вычислений, критическое значение критерия для уровня значимости $\alpha = 0,05$, с учетом степеней свободы равных 31 и 32, $\varphi_{\text{набл}} = 1,4 < \varphi_{\text{кр}} < 2$. Следовательно, нулевая гипотеза о равенстве генеральных дисперсий на уровне значимости 0,05 принимается.

Вывод: Так как наблюдаемое (эмпирическое) значение критерия Фишера меньше критического (табличное) значения, то статистически значимых различий дисперсий в первой и второй группах нет и, следовательно, уровень стрессоустойчивости у учителей и менеджеров одинаков.

Вопросы и задания для самоконтроля

1. *Дайте определения: нормальное распределение, нормальная площадь, нормальное отклонение, колоколообразное распределение.*

2. *Объясните, в каком смысле «нормальная» кривая является нормальной.*

3. *Что такое «сигма-единица», «сигма-точка»?*

4. *Каковы виды идеальных распределений?*

5. *Опишите, как распределения делятся в зависимости от их симметрии? Что могут означать подобные отклонения от «идеальных» распределений?*

6. *Установите долю случаев, лежащих между каждой парой сигма-точек на основной линии нормальной кривой. Представьте результат графически.*

0,3 до 1,6	1,1 до 1,2	−2,58 до +2,58
0,3 до −1,6	0,0 до 1,0	1,5 до 3,0
0,1 до 0,2	1,0 до 2,0	−2,3 до +2,3

7. *Между какими двумя сигма-точками на основной линии нормальной кривой лежат средние 50% случаев?*

8. Объясните, почему доля случаев между 0 и 1,0 сигма не равна доле между 1,0 и 2,0?

9. Среднее нормальное распределение равно 75, а стандартное отклонение равно 3: А) какая доля значений лежит между 72 и 78? Б) какие значения находятся в верхнем дециле распределения? В) какая приблизительно доля лежит между 69 и 81?

10. Перечислите критерии нормальности распределения. Каков смысл этих критериев с точки зрения эмпирической социологии?

11. В чем разница между двусторонними и односторонними статистическими критериями?

12. Что понимается под мощностью статистического критерия?

13. «Нулевая гипотеза»: определение понятия, пример формулировки.

14. Чем статистическая гипотеза отличается, например, от той, которая формулируется социологом в программе эмпирического исследования?

15. Каковы типы критериев значимости? В чем между ними разница?

16. Определите сущность и опишите принцип нахождения критерия асимметрии (A_s).

17. В чем особенности критерия Пирсона Хи-квадрат (как критерия проверки гипотезы H_0)?

18. Что такое «число степеней свободы» и зачем обязательно определять это число при проверке статистического критерия?

19. В чем заключается сущность Т-критерия Стьюдента (как критерия, основанного на нормальном распределении)?

20. Чем отличаются параметрические и непараметрические критерии?

21. В каких случаях исследователя интересует F-критерий Фишера?

Раздел IV. ИЗМЕРЕНИЕ СВЯЗИ МЕЖДУ ПРИЗНАКАМИ С ПОМОЩЬЮ МАТЕМАТИЧЕСКИХ МЕТОДОВ

4.1. Понятие статистической связи

Принцип взаимной сопряженности. Аксиомой науки, а также здравого смысла является убеждение в том, что ни одно событие в природе не возникает «вдруг», но всегда при вполне определенных, известных или неизвестных обстоятельствах. Событие, таким образом, никогда не следует рассматривать изолированно. Оно должно рассматриваться как результат совместных воздействий многих сил, каждая из которых влияет на наблюдаемый результат. Размер семьи, например, может зависеть от таких факторов, как продолжительность супружеской жизни, уровень дохода, степень занятости матери на работе вне дома, ее религиозные взгляды и т. д. Одни из этих факторов имеют тенденцию увеличивать, другие – уменьшать интенсивность изменения и даже препятствовать возникновению события. Так, религиозный фактор может благоприятствовать росту семьи, тогда как экономический фактор может притормозить эту тенденцию. Во всяком случае, нельзя предсказать размер семьи абсолютно точно, можно лишь связывать его с некоторыми факторами или переменными, с которыми сопряжен результат. Это знание основано, таким образом, на принципе взаимной сопряженности. Именно в соответствии с этим принципом и естественные, и социальные, и любые другие науки устанавливают свои методы и цели: 1) идентифицировать переменные (факторы), связанные с событием; 2) раскрыть способы взаимосвязи между факторами и событием; 3) измерить силу этой взаимосвязи.

Первая проблема заключается в выделении сопряженных факторов. Существует бесчисленное количество факторов, отношения между которыми столь запутаны, что «конечный»

фактор при появлении события никогда не может быть получен и тем более измерен. Однако в бесконечных поисках определенности здравый смысл начинает конструировать образцы зависимостей, на базе которых он стремится охватить прошлое, понять настоящее и тем самым предвидеть будущее.

Виды зависимости. *Корреляция (корреляционная зависимость)* – статистическая взаимосвязь двух или нескольких случайных величин (либо величин, которые можно с некоторой допустимой степенью точности считать таковыми). При этом изменения значений одной или нескольких из этих величин сопутствуют систематическому изменению значений другой или других величин. Виды подобной взаимосвязи, конечно, рассматриваются в ходе каждодневного опыта, накапливаемого методом проб и ошибок. В этом процессе наблюдатель мысленно квантифицирует и суммирует свои наблюдения, во-первых, отмечая факторы, которые, как ему кажется, «производят» событие или просто связаны с ним, и, во-вторых, отмечая частоту, с которой он успешно предвидит или предсказывает. Он интуитивно использует принцип корреляции.

Корреляционные связи различаются по форме, направлению и степени (силе) [6; 9; 14].

По форме корреляционная связь может быть прямолинейной или криволинейной. *Прямолинейной* может быть, например, связь между количеством посещенных занятий по курсу «Математические методы в социологии» во время семестра и количеством правильно решенных задач на экзамене по этому курсу. *Криволинейной* может быть, например, связь между уровнем мотивации и правильностью решения задач: при повышении мотивации (гарантированный, на определенных условиях «автомат» по экзамену) эффективность (правильность) выполнения задач сначала возрастает. Оптимальному уровню мотивации соответствует максимальная эффективность выполнения задач. Вполне возможно, что дальнейшему повышению мотивации может сопутствовать снижение эффективности.

По направлению корреляционная связь может быть положительной («прямой») и отрицательной («обратной»). При *положи-*

тельной прямолинейной корреляции более высоким значениям одного признака соответствуют более высокие значения другого, а более низким значениям одного признака – низкие значения другого (см. Рис. 4.1). При отрицательной корреляции соотношения обратные. Если воспользоваться одним из уже рассматриваемых нами примеров, то в случае прямой связи мы имеем: «с возрастанием количества посещений занятий по курсу «Математические методы в социологии» эффективность решения задач возрастает. В случае же обратной связи: «с возрастанием количества посещений занятий по курсу «Математические методы в социологии» эффективность решения задач уменьшается» (что, конечно же, невероятно, и может рассматриваться только для иллюстрации направления связи!).

При положительной корреляции коэффициент корреляции имеет положительный знак, например $r = +0,207$, при отрицательной корреляции – отрицательный знак, например $r = -0,207$ [14, с. 44].

Степень, сила или теснота корреляционной связи определяется по величине коэффициента корреляции.

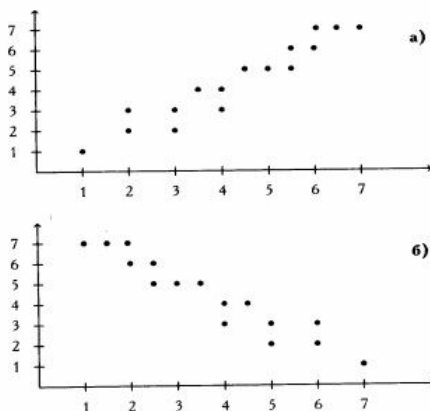


Рис. 4.1. Схема прямолинейных корреляционных связей:
а) положительная (прямая) связь; б) отрицательная (обратная) связь

Сила связи не зависит от ее направленности и определяется по абсолютному значению коэффициента корреляции. Коэффициент корреляции – это величина, которая может варьировать в пределах от -1 до $+1$. В случае полной положительной корреляции этот коэффициент равен « $+1$ », а при полной отрицательной – « -1 ».

На рисунке 4.2 представлено множество распределений двух признаков (X и Y) с соответствующими коэффициентами корреляций между этими признаками для каждого из распределений. Коэффициент корреляции отражает «зашумлённость» линейной зависимости (верхняя строка), но не описывает наклон линейной зависимости (средняя строка), и совсем не подходит для описания сложных, нелинейных зависимостей (нижняя строка). Для распределения, показанного в центре рисунка, коэффициент корреляции не определен, так как его дисперсия равна нулю.

Метод обработки статистических данных, с помощью которого измеряется теснота связи между двумя или более переменными, носит название «*корреляционный анализ*». Корреляционный анализ тесно связан с *регрессионным анализом* (также часто можно встретить термин «*корреляционно-регрессионный анализ*»), который является более общим статистическим понятием [10]. С помощью регрессионного анализа определяют

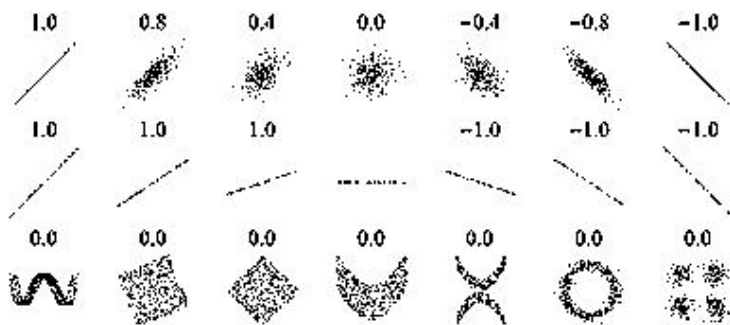


Рис. 4.2. Визуализация двумерных распределений по признакам X и Y , иллюстрирующая силу и направление связи между этими признаками

необходимость включения тех или иных факторов в уравнение множественной регрессии, а также оценивают полученное уравнение регрессии на соответствие выявленным связям, используя *коэффициент детерминации*. В целом, задачи регрессионного анализа лежат в сфере установления формы зависимости, определения функции регрессии. А задачи корреляционного анализа сводятся к измерению тесноты связи между признаками. Величину, силу и направление связи между признаками показывают коэффициенты корреляции. Коэффициент корреляции определяется с помощью нахождения ковариации. *Ковариация* (или *корреляционный момент*) – числовая характеристика совместного распределения двух случайных величин равная математическому ожиданию произведения отклонений случайных величин от их математических ожиданий.

Для некоторых целей грубая субъективная оценка корреляции является вполне достаточной, но для научных целей желательны более точные измерения. Трудность заключается в сложности и разнообразии связей между социальными явлениями. Некоторые черты этой сложности могут быть сформулированы следующим образом: 1) каждое событие является результатом действия многочисленных факторов; 2) сила воздействия факторов изменяется по интенсивности; 3) она может быть направлена в одном или нескольких направлениях; 4) факторы находятся в постоянном взаимодействии; 5) они могут усиливать, противодействовать и уничтожать влияние друг друга.

Эта проблема, вероятно, менее остра в физических науках, чем в науках социальных. Физик посредством лабораторного контроля до определенной степени в состоянии выделять свой объект и манипулировать им. Он может тщательно повторить свои наблюдения в прежних условиях, в то время как социолог часто вынужден пользоваться данными, подобными необработанной руде, и собирать материал из разбросанных источников. Социолог, таким образом, вынужден использовать статистический контроль, так как лабораторный контроль ему недоступен.

Что является доказательством наличия связи? Как можно определить, что факторы находятся во взаимной связи? А, обна-

ружив связь, как можно определить степень или интенсивность этой связи? Вообще говоря, существуют два отличительных признака такого рода связи:

- 1) совместное появление качественных признаков;
- 2) параллельное изменение в двух или более рядах количественных переменных.

Во-первых, относительная частота, с которой появляются вместе определенные качественные признаки, может служить наиболее простым основанием для заключения об ассоциации (связи). В этом и заключается принцип совместного появления признаков. Статистические переменные, подобно людям, оцениваются обычно «в зависимости от общества, в котором они находятся». Например, если преступность более часто обнаруживается среди юношей, чем среди девушек, то заключают, что преступность ассоциируется с признаком «быть юношей». Сила этой связи будет изменяться в зависимости от других факторов, таких, как возраст юношей, характер преступления и многих других признаков. И все они будут затруднять статистическое применение этого простого, на первый взгляд, принципа. Следовательно, едва ли нужно еще раз повторять, что некоторая система группирования и классификации необходима не только как средство установления наличия связи, но и как средство определения силы этой связи.

Во-вторых, если для двух рядов количественных данных изменение в одном из них соответствует с некоторой степенью устойчивости сравнимому изменению в другом ряду, то заключают, что они как-то связаны между собой и что существует связь между двумя рядами данных. Например, по мере того как падает доход, намечается тенденция к увеличению размера семьи; и если продолжить наблюдения на протяжении достаточно обширного ряда, то есть рассмотреть целый ряд семей различных размеров с различным доходом, то очевидность связи усиливается. Этот подход к оценке связи был назван принципом ковариации.

Способы и принципы измерения связи. Интерпретация коэффициентов корреляции. Техника измерения связи должна соответствовать: 1) природе данных; 2) числу взаимодействующих

щих переменных; 3) видам зависимости между ними. Поэтому способы измерения связи будут различаться в зависимости от того, будут ли данные представлены в форме качественных признаков, которые просто перечислены, или в виде количественных измерений, и еще будет ли вид зависимости между переменными простым или сложным.

Используется две системы классификации корреляционных связей по их силе: общая и частная [14, с. 44–45].

Общая классификация корреляционных связей:

1) сильная, или тесная при коэффициенте корреляции $\text{Коэф} > 0,70$;

2) средняя при $0,50 < \text{Коэф} < 0,69$;

3) умеренная при $0,30 < \text{Коэф} < 0,49$;

4) слабая при $0,20 < \text{Коэф} < 0,29$;

5) очень слабая при $\text{Коэф} < 0,19$.

Для качественной оценки тесноты связи часто используют так называемую шкалу Чеддока, представленную в таблице 4.1.

Таблица 4.1

Оценка тесноты связи по шкале Чеддока

Количественная мера тесноты связи	Качественная характеристика силы связи
(0,1–0,3)	Слабая
[0,3–0,5)	Умеренная
[0,5–0,7)	Заметная
[0,7–0,9)	Высокая
(0,9–0,99)	Весьма высокая

Частная классификация корреляционных связей:

1) высокая значимая корреляция при r , соответствующем уровню статистической значимости $p \leq 0,01$;

2) значимая корреляция при r , соответствующем уровню статистической значимости $p \leq 0,05$;

3) тенденция достоверной связи при r , соответствующем уровню статистической значимости $p \leq 0,10$;

4) незначимая корреляция при r , не достигающем уровня статистической значимости.

Две эти классификации, – общая и частная, не совпадают. Первая ориентирована только на величину коэффициента корреляции, а вторая определяет, какого уровня значимости достигает данная величина коэффициента корреляции при данном объеме выборки. Чем больше объем выборки, тем меньшей величины коэффициента корреляции оказывается достаточно, чтобы корреляция была признана достоверной. В результате при малом объеме выборки может оказаться так, что сильная корреляция окажется недостоверной. В то же время при больших объемах выборки даже слабая корреляция может оказаться достоверной [14, с. 45].

Обычно принято ориентироваться на вторую классификацию, поскольку она учитывает объем выборки. Вместе с тем необходимо помнить, что сильная, или высокая, корреляция – это корреляция с коэффициентом $r > 0,70$, а не просто корреляция высокого уровня значимости.

Из многочисленных способов измерения связи рассмотрим только те из них, которые довольно просты в вычислении и наиболее часто используются в социологической практике.

Формулы измерения связи могут быть удобно сгруппированы на основе двух рассмотренных выше принципов связи:

1) *совместного появления* (коэффициент взаимной сопряженности Пирсона (C), стандартные коэффициенты связи Чупрова и Крамера (T и T_c);

2) *ковариации* (метод корреляции ранговых различий Спирмена (r_s) и Кенделла (τ), коэффициент корреляции Пирсона (r), корреляционное отношение (η), коэффициенты множественной ($R_{y(1...k)}$) и частной корреляции ($r_{y1.2}$).

Все формулы и алгоритмы вычисления этих коэффициентов будут в подробностях и с применением реальных примеров рассмотрены нами, в подразделах 4.4–4.7 данного издания.

Коэффициенты рассчитываются с использованием стандартных программ и показывают меру взаимообусловленности в распределении частот появления соответствующих признаков. Один из признаков условно считается зависимым, другой –

детерминирующим, однако заключение о наличии связи может дать только качественный анализ всей совокупности связей. Анализ коэффициентов связи позволяет:

- выделить факторы, статистический уровень влияния которых позволяет исключить их из дальнейшего анализа (гипотеза о наличии связи отрицается);
- проранжировать оставшиеся связи по уровню взаимной сопряженности с изучаемым процессом, при этом следует иметь в виду, что уровень взаимной сопряженности может определяться как влиянием данного фактора на процесс, так и взаимным изменением данного фактора и процесса под влиянием третьего фактора.

В заключение данного параграфа, приступая к более подробному и конкретному рассмотрению процедуры измерения взаимосвязей между признаками, хотелось бы обратить внимание на два следующих момента:

!!! во-первых, исходя из логических рассуждений, говоря о корреляции, не стоит отождествлять понятия «связь» и «зависимость». Если между признаками есть корреляционная связь – это еще не значит, что они взаимозависимы. Наличие связи между признаками может свидетельствовать не о зависимости этих признаков между собой, а о зависимости обоих этих признаков от какого-то третьего признака или сочетания признаков, не рассматриваемых в исследовании. Зависимость подразумевает влияние – любые согласованные изменения, которые могут объясняться сотнями причин [14, с. 40–41];

!!! во-вторых, как удачно отмечает автор одного из российских учебников по применению математических методов в психологии, Л. С. Титкова, корреляционные связи не могут рассматриваться как свидетельство причинно-следственной связи. Они свидетельствуют лишь о том, что изменениям одного признака, как правило, сопутствуют определенные изменения другого, но находится ли причина изменений в одном из признаков или она оказывается за пределами исследуемой пары признаков, нам неизвестно [14, с. 40–41].

4.2. Корреляционное поле

Корреляционная зависимость — связь между признаками, состоящая в том, что в зависимости от применения одного признака меняется величина другого. Как правило, при изучении взаимозависимости двух признаков различают: независимые признаки (факторные), которые чаще всего обозначаются — X ; зависимый признак (результатирующий) — Y .

В ходе корреляционного анализа необходимо узнать, как под влиянием факторных признаков изменяется результирующий, если он изменяется вообще и по какому-либо закону. Вспомогательным средством анализа выборочных данных является корреляционное поле. Если даны значения двух признаков X и Y , эти значения можно сопоставить путем отражения их в системе координат и нанесении на плоскость точек, соответствующих этим значениям и, при необходимости, соединения этих точек непрерывной линией. Расположение точек позволяет сделать предварительное заключение о характере и форме зависимости. Корреляционное поле относится к двумерной совокупности данных точно так же, как гистограмма — к одномерной совокупности. Оно наглядно изображает распределение наблюдений в целом, тем самым позволяя получить грубую, но полезную оценку степени корреляции, прежде чем вычислить последнюю. Этот необходимый прием уже был использован для того, чтобы изобразить тенденцию временных рядов, однако детали построения поля еще не были описаны.

Чтобы проиллюстрировать более полно принципы построения корреляционного поля и его использования, возьмем в качестве примера таблицу 4.2, где представлена динамика развития сети высших учебных заведений Украины III–IV уровня аккредитации и численности студентов в них за 1992/2000 годы.

Для того чтобы свести данные этой таблицы к корреляционному полю, сначала чертят горизонтальную и вертикальную оси, как и при построении гистограммы. Оси задаются как приблизительно равные по длине, если только нет достаточного основания для отступления от этого правила. Точно так же по

Таблица 4.2

**Динамика развития сети высших учебных заведений Украины
III–IV уровня аккредитации и численности студентов в них**

<i>Годы</i>	<i>Количество высших учебных заведений Украины III–IV уровня аккредитации</i>	<i>Численность студентов в них (тыс.)</i>
[1992 – 1993)	158	718,8
[1993 – 1994)	159	680,7
[1994 – 1995)	232	645
[1995 – 1996)	255	617,7
[1996 – 1997)	274	595
[1997 – 1998)	280	526,4
[1998 – 1999)	298	503,7
[1999 – 2000)	313	503,7

обыкновенно независимая переменная располагается вдоль горизонтальной оси, а зависимая переменная – вдоль вертикальной.

Затем на осях устанавливаются шкалы таким образом, чтобы согласовать наблюдаемые области значений соответствующих переменных. Так, горизонтальная шкала охватывает промежуток от 100 до 350 единиц – количество высших учебных заведений, в то время как вертикальная шкала простирается от 500 до 800 – численность студентов. Нет необходимости говорить о том, что каждую ось следует разбить на достаточное число делений, чтобы обеспечить точное и сравнительно простое нанесение точек. В отличие от гистограммы, вертикальная шкала на поле не обязательно должна начинаться с нуля по той причине, что основное внимание здесь сосредоточено на очертании рассеивания, а не на относительной частоте, которая оценивается по высоте кривой линии.

Начертив оси и прошкаливовав их, можно изобразить каждую пару значений в виде точки на плоскости. Для каждой пары значений переменных, обозначенных точкой, значение Y определяет высоту точки над горизонтальной линией, а значение X определяет ее удаленность от вертикальной оси. Так, 1996–97 учебный год изображен на рисунке 4.3 точкой, расположенной на пересечении направляющих линий, восстановленных перпендикулярно

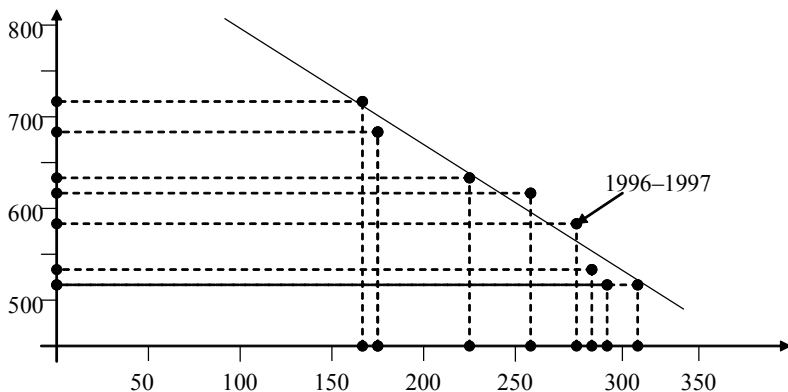


Рис. 4.3. Корреляционное поле динамики сети высших учебных заведений Украины III–IV уровня аккредитации и численности студентов в них

к соответствующим осям из точек $Y = 595$ и $X = 274$. Для всех остальных пар чисел находятся «свои» точки в заданной системе координат, установленные на пересечении взаимно перпендикулярных прямых. Совокупность всех таких точек образует корреляционное поле.

Виды рассеивания. Именно структура рассеивания позволяет судить о характере связи, а такое заключение является важной предпосылкой точного измерения связи. Так, например, как это явствует из анализа корреляционного поля, постепенное увеличение количества высших учебных заведений III–IV уровня аккредитации в Украине с 1992 по 2000 год сопровождается постепенным снижением численности студентов в них, то есть численность студентов изменяется на постоянную величину при изменении количества вузов на единицу измерения. Такая зависимость называется прямолинейной, или просто линейной, потому что общее направление распределения точек близко к прямой линии.

Любая такая линия концентрации данных, проведенная на глазок от руки или построенная на строгой математической основе, называется линией регрессии. Это понятие было введено

в обращение Ф. Гальтоном в 1877 году, который использовал его в связи с корреляционным исследованием некоторых особенностей родителей и их детей. Он заметил, что такая линия эффективно выражает тенденцию среди детей «регрессировать» к среднему уровню родителей по целому ряду признаков. Термин сохранился и получил широкое распространение, хотя его значение несколько изменилось.

Зависимость, изображенную на рисунке 4.4, называют обратной, так как два ряда значений переменных движутся в противоположных направлениях: по мере увеличения количества вузов численность студентов, обучающихся в них, падает. Если бы численность студентов и количество вузов возрастали одновременно, то общее направление рассеяния точек было бы снизу вверх и слева направо. Такая линия регрессии, как на рисунке 4.4, является доказательством прямой линейной зависимости между двумя переменными.

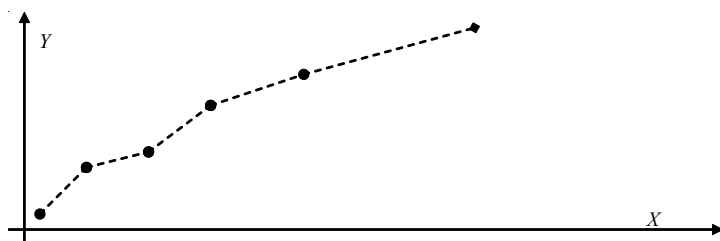


Рис. 4.4. Пример прямой ковариации

Тенденция рассеяния не всегда бывает линейной, чаще всего она бывает криволинейной и принимает форму одного из многочисленных видов кривых. На рисунке 4.5 дается пример, который изображает зависимость между благосклонностью отношения к национальному меньшинству и интенсивностью этого отношения. Решительное мнение – «За» или «Против» сочетается со значительной определенностью отношения и наоборот.

Все три рассмотренных корреляционных поля имеют ярко выраженную тенденцию в рассеянии точек, поэтому каждый мог бы без особого труда охарактеризовать эти тенденции как линейные или криволинейные. Но распределения эмпирических

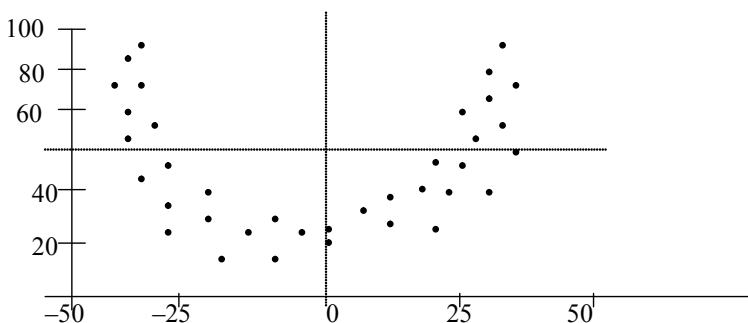


Рис. 4.5. Корреляционное поле (гипотетические данные)

наблюдений редко бывают столь определенными и недвусмысленными; более часто обе эти тенденции сочетаются в одной и той же совокупности данных и тем самым усложняют проблему выражения корреляции в виде обобщенной меры. Например, зависимость, представленная на рисунке 4.5, может быть интерпретирована как прямолинейная, однако, очевидно, что кривая линия больше соответствует тенденции рассеяния. По мере того как признак увеличивается, признак также увеличивается, но в постепенно замедляющемся темпе. Это подтверждается выравниванием в скоплении точек. Наблюдение на корреляционном поле отклонений от общей кривой требует специального анализа, так как они представляют собой нарушение «закона» связи. Однако в основном закон связи выдерживается довольно хорошо, давая меру предсказуемости изменений одной переменной в зависимости от изменений другой. Если знаем, например, что количество высших учебных заведений III–IV уровня аккредитации Украины составляет 200, то можем предсказать, что численность студентов в них примерно равна 650 тыс., что соответствует высоте линии регрессии в этой точке (см. Рис. 4.3).

Несомненно, такое предсказание было бы не свободно от ошибки по той очевидной причине, что ни одно из наблюдаемых значений не попадает прямо на кривую в этой точке – все отклоняются в большей или меньшей степени. Очевидно, точность любого такого предсказания изменялась бы в зависимости от

того, насколько плотно прилегают точки к линии, выражающей связь между двумя рядами данных. Точность предсказания, а следовательно, и корреляция, была бы высокой, если бы точки располагались концентрированно; когда точки рассеяны очень широко, точность предсказания бывает соответственно невысокой. Предсказание и корреляция были бы абсолютно полными только тогда, когда все точки располагались бы на линии регрессии. При другой крайности, когда рассеяние точек носит совершенно случайный характер, можно с равным успехом игнорировать так называемую «переменную-индикатор».

Скедастичность (вариабельность). Зная значение одного признака, не всегда можно с одинаковой вероятностью предсказать изменение другого, это связано со степенью рассеивания признака. Так, для примера, показывающего динамику сети высших учебных заведений Украины III–IV уровня аккредитации и численности студентов в них (см. *Табл. 4.2*), можно с одинаковой вероятностью предсказать изменение одного признака по изменению другого, так как рассеивание признака – количества вузов равно $313 - 158 = 155$, а рассеивание признака – численности студентов в этих вузах равно $718,8 - 503,7 = 215,1$ тыс., что составляет числа одного порядка.

Рассеяние значений Y , соответствующих данному значению X , называется **скедастичностью**. Если степень вариации значений (ширина зоны рассеяния) одинакова для всех значений X , то можно говорить о том, что переменная Y *гомоскедастична* по отношению к X . В противном случае, если, например, степень рассеяния значений Y уменьшается по мере изменения X , говорят, что переменная Y *гетероскедастична* по отношению к X . Гетероскедастичность означает, что степень корреляции неодинакова для всей совокупности данных, следовательно, ее наличие уменьшает возможность обойтись одним обобщенным показателем корреляции, который, в конце концов, всегда является средней величиной. Подобно тому, как не всегда можно вычислить среднее арифметическое для гетерогенных данных, точно так же избегают рассчитывать среднюю меру корреляции для гетероскедастического рассеяния.

Гетероскедастичность может быть ярко выраженной. Так, рассеяние может быть грушевидным, гантелеобразным или *j*-образным. Эти странного вида диаграммы рассеяния никоим образом не исчерпывают всех возможных типов распределения, которые могут встречаться в практической работе. Однако они все же пригодны для того, чтобы подтвердить полезность такого рода визуальных средств. Хотя диаграмма рассеяния не дает математической меры корреляции, она все же указывает: а) является ли зависимость простой прямолинейной или более сложной, б) является ли зависимость устойчивой для всех значений переменной, в) является ли связь сильной или слабой. Диаграмма рассеивания является необходимым средством анализа и играет при изучении ковариации ту же роль, что и график распределения частот при обработке одномерной совокупности данных. Она позволяет обзирать все распределение в целом.

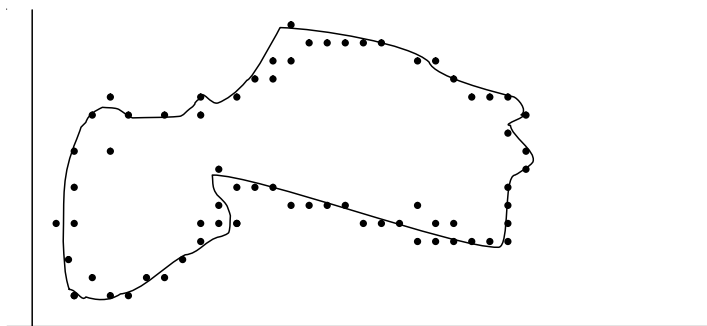


Рис. 4.6. Диаграмма рассеяния

4.3. Корреляционная таблица

Вместо того чтобы наносить на чертеж отдельные наблюдения, можно сгруппировать их. Такая группировка отвечает следующим целям: 1) определению основного вида рассеивания и 2) облегчению статистической обработки данных. При группировании двумерных совокупностей данные классифицируются

так, что каждый случай учитывается одновременно по двуразрядным интервалам, тем самым располагая каждый случай в клетке на пересечении данной строки и данного столбца. Следовательно, получается так, как если бы наложили координатную сетку на корреляционное поле, подсчитали точки в клетках и вписали в каждую клетку соответствующее число. В результате применения такой операции и такого распределения получают таблицу распределения совместных частот (см. *Табл. 4.3*).

Естественно, в любой реальной ситуации не стали бы действовать таким несколько странным образом, скорее всего, непосредственно разнесли бы неупорядоченные двумерные наблюдения по клеткам координатной сети, специально разработанной для этой цели, и затем подсчитали бы количество наблюдений в каждой клетке.

Техника группирования. Чтобы уяснить технику группирования, перечислим правила, которых следует придерживаться:

1. Подбирают подходящий интервал группировки для каждой переменной.

2. Наносят эти интервалы на соответствующие оси координат и из каждой отметки проводят направляющие линии таким образом, чтобы получить координатную сетку. Предполагается, что каждая клетка представляет собой место пересечения двух интервалов группировки.

3. Помещают каждую связанную пару значений в соответствующую клетку и обозначают ее наличие каким-либо значком.

4. Подсчитывают значки в каждой клетке и заменяют их числом. Каждое такое число представляет собой совместную частоту, т. е. частоту, с которой появляются вместе значения или точки, относящиеся к двум признакам.

5. Суммируют по строкам и столбцам для того, чтобы получить маргинальные суммы. Они дают простые частотные распределения каждой переменной.

6. Суммируют маргинальные частоты, чтобы получить общее число всех случаев – N , причем, сумма частот по столбцам является контрольной цифрой для суммы частот по строкам и наоборот.

Таблица 4.3

Корреляционная таблица совместных частот распределения времени (Y) в часах в неделю, которое тратят респонденты на занятия, не связанные с учебой в школе, в зависимости от класса обучения (X)

X\Y	до 1	1–2	2–3	3–4	4–5	5–6	6–7	7–8	8–9	9–10	Σ
1	6										6
2	1	6									7
3	2	4	8	2		4					20
4			6	7	2	3					18
5			5	3	4	3					15
6				2	1	4	3				10
7						5	3	2	3		13
8							6	2			8
9								1	2	1	4
10								1	1	2	4
11							1	2	2	3	8
12									4	1	5
Σ (маргинальные суммы по столбцам)	9	10	19	14	7	19	13	8	12	7	N = 118

Таким образом, построение корреляционной таблицы происходит в соответствии с четко сформулированными принципами: классификация должна быть достаточно детальной, чтобы обнаружить форму распределения совместных численностей, и в то же время не настолько детальной, чтобы некоторые ряды и столбцы не остались совершенно пустыми. Но поскольку вначале неизбежны ошибки, то, вероятно, желательно иметь скорее избыток клеток в таблице, чем недостаток. Обычно всегда легче объединить клеточные частоты, чем разделить их.

Функция корреляционной таблицы. Необходимо признать, что корреляционная таблица не может изобразить вид зависимости так наглядно, как корреляционное поле (коррелограмма) по той причине, что цифры не могут передать столь же эффективно оттенки в плотности распределения, как это делает

рассеяние точек. Тем не менее, несмотря на относительную грубость корреляционной таблицы, она часто оказывается вполне пригодной для установления вида зависимости и правильного ее понимания. Кроме того, распределения в строках и столбцах корреляционной таблицы позволяют произвести статистическое измерение скедастичности, что невозможно было бы сделать на основании исходного корреляционного поля. Для этого достаточно только вычислить квадратическое отклонение для каждого ряда или столбца. Близкое сходство между квадратическими отклонениями служило бы доказательством гомоскедастичности, тогда как заметные различия свидетельствовали бы о гетероскедастичности.

К тому же маргинальные распределения переменных X и Y , которые, безусловно, являются неотъемлемыми чертами корреляционной таблицы, также являются важным и общепризнанным средством для определения возможной степени связи между переменными. В частности, маргинальные распределения устанавливают границы степени изменчивости получаемой корреляции. Например, непохожие маргинальные распределения исключают полную линейную зависимость. Обычно маргинальные распределения всегда действуют тем или иным образом в направлении ограничения силы связи между спаренными переменными; следовательно, всегда необходимо изучить их для того, чтобы определить ограничения, которые они накладывают, и возможности для обоснованного применения данного корреляционного показателя.

Корреляционная таблица может служить в качестве вспомогательного средства вычисления в тех случаях, когда число наблюдений велико или когда не имеется вычислительных машин. Наконец, она является предпосылкой для очень простого перехода к криволинейной корреляции, как будет показано в следующем параграфе.

В данном параграфе подчеркнута важная роль поля корреляции и корреляционной таблицы в предварительном анализе ковариации. С помощью этих диаграмм можно определить, является ли зависимость: а) линейной или криволинейной; б) прямой

или обратной; в) слабой или сильной; г) являются ли переменные гомоскедастичными по отношению друг к другу; д) имеются ли значительные отклонения от основной тенденции или «закона связи» и е) являются ли маргинальные распределения симметричными или скошенными и сопоставимы ли они. Имея в виду получение этих важных предварительных сведений, всегда необходимо еще до измерения корреляции построить и тщательно изучить» диаграмму рассеяния или корреляционную таблицу, или даже и то и другое одновременно. Можно использовать неподходящие методы и, следовательно, прийти к неправильным выводам, если не прибегать к помощи этих визуальных средств.

Может это и покажется повторением, но учитывая то, что с началом следующего подраздела мы переходим к детальному рассмотрению коэффициентов связи, все-таки, напомним, что формулы измерения связи могут быть удобно сгруппированы на основе двух рассмотренных выше принципов связи: 1) *совместного появления* (коэффициент взаимной сопряженности Пирсона (C), стандартные коэффициенты связи Чупрова (T) и Крамера (T_c), коэффициент ассоциации Юла (Q) и коэффициент контингенции Пирсона (Φ); 2) *ковариации, взаимоизменяемости* (метод корреляции ранговых различий Спирмена (r_s) и Кенделла (τ), коэффициент корреляции Пирсона (r), коэффициент λ Гуттмана, корреляционное отношение (η), коэффициенты множественной ($R_{y(1...k)}^2$) и частной корреляции ($r_{y1.2}$).

4.4. Коэффициенты связи, основанные на Хи-квадрат Пирсона (для номинальных признаков)

Номинальные признаки – это те признаки, которые измерены с помощью номинальных шкал. В первых параграфах данного издания мы говорили о том, что номинальные шкалы – шкалы низкого типа. А это значит, что возможное число измерительных процедур, которые можно осуществить применительно к таким шкалам, крайне ограничено. Считается, что для расширения исследовательских возможностей желательно обращаться именно

к шкалам более высокого типа (порядковым, метрическим). Возможно, это, действительно так, но не в социологии. Согласимся с мнением тех социологов, которые утверждают, что роль номинальных данных в социологии огромна. В частности, Ю. Н. Толстова¹⁷ объясняет это утверждение следующим образом [15, с. 164–165]:

- во-первых, именно номинальные данные чаще всего используются социологами, что объясняется сравнительной простотой их получения, естественностью интерпретации;
- во-вторых, номинальные данные являются более надежными, чем данные, полученные по шкалам более высокого типа, в том смысле, что за ними обычно не стоят трудно проверяемые модели восприятия (имеется в виду восприятие респондентом предлагаемых ему для оценки объектов, суждений, мнений и т. д.

Некоторые же авторы, например, С. В. Чесноков¹⁸, вообще полагают, что в социологии только номинальные шкалы имеют право на существование [17]. Нередки случаи, когда количественные признаки, для которых более естественным кажется измерение с помощью метрических (числовых) шкал, вполне успешно измеряются с помощью номинальных шкал. Несмотря на то что номинальные шкалы являются шкалами низкого типа, для их анализа имеется немало эффективных методов. Наиболее часто социологи прибегают к расчетам коэффициентов, основанных на Хи-квадрат, то есть связанных с обязательным определением

¹⁷ **Толстова Юлианна Николаевна**, д-р социол. наук, профессор кафедры сбора и анализа социологической информации Высшей школы экономики (г. Москва, Россия).

¹⁸ **Сергей Валерианович Чесноков** (р. 29 июня 1943 г., СССР) – российский ученый, математик, социолог, культуролог, музыкант, специалист по методам анализа данных и применению математических методов в гуманитарных исследованиях и проектах. Известен как создатель детерминационного анализа и детерминационной логики, исследователь гуманитарных оснований точных наук, активный участник песенного движения и артистического андеграунда в СССР и современной России.

меры различия между наблюдаемыми (эмпирическими) и теоретическими частотами. Сам по себе Хи-квадрат является свидетельством связи между двумя признаками. Понятно, что при отсутствии связи величина Хи-квадрат равна нулю, и это значение является минимальным. В подразделе 3.4 данного издания мы приводили достаточно развернутый пример вычисления χ^2 , в связи с чем не считаем необходимым обращаться к подобным примерам вновь. Отметим, что основная проблема заключается в невозможности определения максимального значения χ^2 , которое могло бы свидетельствовать о силе связи. Эта величина не имеет общего для всех таблиц сопряженности максимального значения, даже тогда, когда связь между признаками является максимально сильной (то есть когда каждому значению (категории) одного признака в точности соответствует значение другого признака). Кроме того χ^2 зависит от числа степеней свободы, что делает невозможным сравнение между собой значений данной величины для таблиц с разным числом строк и столбцов. Эти аргументы, указывающие на несовершенство χ^2 , приводят к осознанию необходимости поиска коэффициентов, которые имели бы фиксированный максимум в случае максимальной связи и позволяли бы сравнивать между собой разные таблицы.

Среди коэффициентов, удовлетворяющих этим требованиям, социологами наиболее часто используются такие, как коэффициент сопряженности Пирсона (C), коэффициенты Чупрова (T) и Крамера (Tc).

Коэффициент сопряженности Пирсона (C). Коэффициент сопряженности Пирсона (C) основывается на отклонении наблюдаемых частот в клетках таблицы от ожидаемых частот и предполагает, что распределение носит случайный характер [9, с. 78–84]. Эти отклонения как раз и измеряются показателем χ^2 , в соответствии с чем, формула вычисления коэффициента сопряженности Пирсона имеет следующий вид:

$$C = \sqrt{\frac{\chi^2}{\chi^2 + n}},$$

где • χ^2 – вычисленная мера различия между эмпирическими и теоретическими частотами таблицы сопряженности признаков;

• n – число единиц наблюдения (объем выборки).

Эта формула, в которой учтено изменяющееся число наблюдений n , дает нормализованный показатель связи C . Необходимо обратить внимание на тот факт, что если χ^2 велико по сравнению с n , то C будет стремиться к единице, так как числитель и знаменатель фактически будут равны; однако, если χ^2 мало по сравнению с n , то коэффициент C будет также мал и в пределе будет стремиться к нулю. Если $\chi^2 = 0$ (то есть, если нет никакого расхождения между полученными данными и чисто случайным распределением), коэффициент C также будет равен нулю, потому что числитель равен нулю. Если все сказанное компактно представить в виде неравенства, будем иметь следующее: $0 \leq C < 1$.

Как видим, коэффициент сопряженности Пирсона C никогда не достигает единицы. Хотя и является очевидным, что чем ближе значение C к единице, – тем сильнее связь, однако представляется необходимым нахождение $C_{\max}$ как максимального значения коэффициента, являющегося действительным эквивалентом полной корреляции.

$C_{\max}$ вычисляется по следующей формуле:

$$C_{\max} = \sqrt{\frac{\min(r-2; c-1)}{1 + \min(r-1; c-1)}},$$

где • r – число строк таблицы сопряженности;

• c – число столбцов таблицы сопряженности;

• $\min$ – означает, что принимается в расчет минимальная разность: либо от числа строк, либо числа столбцов.

Для квадратных таблиц уместно применить следующую формулу нахождения:

$$C_{\max} = \sqrt{\frac{t-1}{t}},$$

где • t – число строк (либо столбцов – разницы нет).

Рассмотрим все сказанное на реальном примере. Повторимся, что алгоритм вычисления χ^2 описывался нами в третьем разделе

данного издания (в подразделе 3.4), не будем повторяться, а просто подставим в формулу коэффициента сопряженности Пирсона C то значение χ^2 , которое мы получили в том примере, пытаясь определить наличие/отсутствие связи между признаками. Однако, исходя из того, что в данном параграфе мы ведем речь только о номинальных признаках, предположим, что мы искали взаимосвязь между признаками «Профессиональные предпочтения» и «Социально-классовое происхождение», измеренными с помощью номинальных шкал. Итак, предположим, что мы получили $\chi^2 = 25,18$.

$$C = \sqrt{\frac{25,18}{25,18 + 154}} = 0,37 \approx 0,4,$$

$$C_{\max} - \sqrt{\frac{\min(4-1)(3-1)}{1 + \min(4-1)(3-1)}} = \sqrt{\frac{\min 3; 2}{1 + \min 3; 2}} = \sqrt{\frac{2}{1+2}} = \sqrt{0,7} = 0,8.$$

Вывод: учитывая значение $C_{\max}$, исходя из общих принципов интерпретации коэффициентов связи, можно сделать вывод о средней силе связи между признаками.

В заключение разговора о коэффициенте сопряженности Пирсона C еще раз подчеркнем, что $C_{\max}$ является действительным эквивалентом полной корреляции, следовательно, этот показатель может быть использован в качестве стандарта для любой связи, которая меньше, чем полная (при условии, что распределения маргиналов являются идентичными). Отношение $C_{\max}$ и C является приблизительно эквивалентным условной мере связи, принимающей значения от «0» до «1». Чем больше размеры таблицы, тем ближе $C_{\max}$ к «1».

Коэффициенты сопряженности Чупрова (T) и Крамера (T_c).

Уже сложившимся стандартом для измерения связи между признаками в прикладном социологическом исследовании являются еще два коэффициента сопряженности, тоже основанные на χ^2 : это – коэффициенты Чупрова (T) и Крамера (T_c). Считается, что эти коэффициенты более строго оценивают тесноту связи, чем коэффициент сопряженности Пирсона (C).

Коэффициент Чупрова (T) находится с помощью следующей формулы:

$$T = \sqrt{\frac{\chi^2}{n\sqrt{(r-1) \times (c-1)}}},$$

где χ^2 – вычисленная мера различия между эмпирическими и теоретическими частотами таблицы сопряженности признаков;

- n – число единиц наблюдения (объем выборки);
- r – число строк таблицы сопряженности;
- c – число столбцов таблицы сопряженности.

Значение, принимаемое коэффициентом Чупрова, может варьироваться от «0» до «1», соответствующее неравенство имеет следующий вид: $0 \leq T \leq 1$.

Как видим, значение данного коэффициента может быть равным единице, однако только в случае квадратной матрицы (то есть такой, которая имеет одинаковое число строк и столбцов; $r = s$).

В том случае если $r \neq s$, необходимо, как и в случае с коэффициентом сопряженности Пирсона (C), найти действительный эквивалент полной корреляции или $T_{\max}$. Формула вычисления этого значения имеет следующий вид:

$$T_{\max} = \sqrt[4]{\frac{\min(r-1; c-1)}{\max(r-1; c-1)}},$$

где r – число строк таблицы сопряженности;

- c – число столбцов таблицы сопряженности;
- $\min$ – означает, что принимается в расчет минимальная разность: либо от числа строк, либо от числа столбцов;
- $\max$ – означает, что принимается в расчет максимальная разность: либо от числа строк, либо от числа столбцов.

Проиллюстрируем сказанное на примере, который рассматривался выше. С помощью коэффициента Чупрова (T) определим силу связи между признаками «Профессиональные предпочтения» и «Социально-классовое происхождение», для которых $\chi^2 = 25,18$.

$$T = \sqrt{\frac{25,18}{154\sqrt{(4-1) \times (3-1)}}} = \sqrt{\frac{25,18}{377,3}} = \sqrt{0,067} = 0,26 \approx 0,3;$$

$$T_{\max} = \sqrt[4]{\frac{\min(4-1; 3-1)}{\max(4-1; 3-1)}} = \sqrt[4]{\frac{2}{3}} = \sqrt[4]{0,67} \approx 0,9.$$

Вывод: учитывая значение $T_{\max}$, очень приближенное к единице, исходя из общих принципов интерпретации коэффициентов связи, можно сделать вывод об умеренной (однако, ближе к слабой) силе связи между признаками.

Возникает вопрос, в пользу какой связи – слабой или умеренной – сделать окончательный вывод? И без того спорная ситуация, связанная с тем, что предварительно полученный коэффициент сопряженности Пирсона свидетельствует в пользу более сильной связи между признаками, усложняется еще и проблемой округления полученного значения. Дело в том, что значение 0,26 – это еще не умеренная, а слабая связь, однако если его округлить до 0,3 (что имеем полное право сделать), можно получить показатель, свидетельствующий об умеренной связи. Как поступить в такой ситуации? Каков должен быть окончательный вывод? Тут-то как раз и «приходит на помощь» коэффициент Крамера. И неспроста эти два коэффициента, Чупрова и Крамера, обычно упоминаются в паре. Значения этих коэффициентов как бы «подкрепляют» друг друга.

Коэффициент Крамера (T_c)¹⁹ также является мерой связи двух номинальных переменных, основанной на критерии X^2 . Однако, учитывая то, что за основу вычисления, как правило, берется значение коэффициента Чупрова, можно сказать, что T_c является уточняющим по отношению к T . Для нахождения коэффициента Крамера применяется следующая формула:

¹ В некоторых источниках можно встретить использование « V » в качестве условного обозначения для коэффициента Крамера.

$$Tc = \frac{1}{T_{\max}},$$

где • T – значение коэффициента Чупрова;

• $T_{\max}$ – максимально возможное значение коэффициента, являющееся действительным эквивалентом полной корреляции.

Значение, принимаемое коэффициентом Крамера, также может варьироваться от «0» до «1», соответствующее неравенство имеет следующий вид: $0 \leq Tc \leq 1$.

Рассмотрим сказанное на примере и найдем коэффициент Крамера для измерения силы связи между теми же признаками «Профессиональные предпочтения» и «Социально-классовое происхождение», для которых $\chi^2 = 25,18$. В итоге получим:

$$Tc = \frac{0,3}{0,9} = 0,33.$$

Вывод: исходя из общих принципов интерпретации коэффициентов связи, можно сделать вывод об умеренной силе связи между рассматриваемыми признаками.

Как видим, приведенный выше пример как нельзя лучше объясняет, почему одного коэффициента недостаточно для измерения связи между признаками.

4.5. Коэффициенты связи для матриц 2×2 (для номинальных и порядковых признаков)

Для определения статистической связи переменных, измеренных дихотомической шкалой наименований (то есть при помощи номинальной шкалы, содержащей только две альтернативы), используются коэффициенты *ассоциации* (Юла, Q) и *контингенции* (Фи, Φ).

Коэффициент ассоциации Юла (Q). Чтобы иметь единую меру связи для матриц 2×2, английский статистик Д. Юл предложил следующий коэффициент связи, который обозначил « Q » в честь известного ученого XIX века А. Кьютелета (Quetelet)²⁰. Соответствующая формула имеет следующий вид:

$$Q = \frac{ad - bc}{ad + bc},$$

где a, b, c, d являются частотами, распределенными в четырехклеточной таблице (матрице 2×2) следующим образом:

Значение, принимаемое коэффициентом ассоциации, может варьироваться от «-1» (при полной обратной связи) до «1» (при

<i>Признаки</i>	<i>A</i>	<i>He-A</i>
<i>B</i>	a	b
<i>He-B</i>	c	d

полной прямой связи) и быть равным «0» при отсутствии статистической зависимости, а значит, и связи между признаками. Соответствующее неравенство имеет следующий вид: $-1 \leq Q \leq 1$.

При этом, чем ближе значение коэффициента Q к нулю, – тем слабее связь между признаками.

Применим приведенную выше формулу вычисления коэффициента ассоциации Q при обработке данных таблицы 4.4.

Таблица 4.4

Распределение правонарушителей по полу

	<i>Юноши</i>	<i>Девушки</i>	Итого
<i>Правонарушители</i>	20	0	20
<i>He-правонарушители</i>	50	50	80
Итого	50	50	100

Как видим из таблицы, в гипотетической группе из 100 человек имеется 20 правонарушителей, и все они – юноши. В данном случае возможно точное предсказание, то есть существует полная определенность относительно пола правонарушителя. Исходя из

²⁰ **Lambert Adolphe Jacques Quetelet** (22 февр. 1796 – 17 февр. 1874) – бельгийский астроном, математик, статистик, социолог, прославившийся успешным использованием математико-статистических методов в социальных науках.

имеющихся данных, можно сказать, что преступность всецело объясняется признаком пола, так как ни одна из девушек не входит в группу правонарушителей. Однако в социологической практике чаще встречаются не такие очевидные ситуации, что требует дополнительных вычислительных процедур для измерения точной корреляции между признаками. Осуществим такую процедуру с использованием соответствующей формулы для вычисления коэффициента ассоциации Юла (Q) и получим:

$$Q = \frac{20 \cdot 50 - 30 \cdot 0}{20 \cdot 50 + 30 \cdot 0} = \frac{1000 - 0}{1000 + 1} = 1.$$

Такое значение показателя ($Q = 1$) является убедительной мерой полной связи между полом и склонностью к совершению преступления. Очевидно, что коэффициент ассоциации достигает единицы в ситуации, когда $b = 0$ либо $c = 0$. Следовательно, сам коэффициент ассоциации показывает нам не только силу и направление связи. Зная этот коэффициент, мы можем сделать конкретный вывод и относительно характера этой связи, то есть сказать с уверенностью, что именно юноши, а не девушки являются правонарушителями. Если сказать более строго, то коэффициент ассоциации измеряет одностороннюю связь, в то время как коэффициент контингенции (Φ), о котором речь пойдет далее, измеряет двустороннюю связь между признаками и на его основе мы не смогли бы сделать вывод относительно характера этой связи. То есть зная о наличии сильной прямой связи между полом и склонностью к правонарушениям, конкретизировать, каков же пол правонарушителей.

Коэффициент контингенции Φ (Φ). Надо сказать, что статистические свойства Φ аналогичны Q , но в некоторых отношениях они отличаются друг от друга. Подобно Q , этот показатель применим только к дихотомическим таблицам 2×2 , не фиксирующим градаций в значениях признаков, или к непрерывным переменным, которые могут быть более или менее обоснованно дихотомизированы. Точно так же как и Q , Φ измеряет интенсивность связи, выражаемой распределением частот в клетках таблицы. Что касается формул для Φ и Q , то числители

у них идентичны, а знаменатели сконструированы различным образом. Формула вычисления коэффициента контингенции Φ_{ii} имеет следующий вид:

$$\Phi = \frac{ad - bc}{\sqrt{(a+b)(a+c)(b+d)(c+d)}},$$

где a , b , c и d – являются уже известными частотами четырехклеточной таблицы, матрицы 2×2 , а суммы отдельных клеточных частот являются маргиналами.

Значение, принимаемое коэффициентом контингенции, может варьироваться от «-1» (при полной обратной связи) до «1» (при полной прямой связи) и быть равным «0» при отсутствии статистической зависимости, а значит и связи между признаками. Соответствующее неравенство имеет следующий вид: $-1 \leq \Phi \leq 1$.

При этом, чем ближе значение коэффициента Φ к нулю, тем слабее связь между признаками. Однако, если коэффициент ассоциации Юла достигает единицы в ситуации, когда $b = 0$ либо $c = 0$, то коэффициент контингенции Φ_{ii} достигает единицы в иной ситуации, а именно, когда $a = d = 0$ или $b = c = 0$.

Для того чтобы продемонстрировать различия в процессе применения Q и Φ , можно обратиться опять к перекрестной классификации правонарушителей по полу в случае, когда Q был равен единице (см. Табл. 4.4). Подставляя в формулу для Φ те же самые данные, получим значение коэффициента в два раза меньше, а именно:

$$\Phi = \frac{1000 - 0}{\sqrt{20 \cdot 80 \cdot 50 \cdot 50}} = \frac{1000}{2000} = 0,5.$$

Помним, что в этом случае Q равен единице, так как все правонарушители являются мальчиками. Однако обратное утверждение, что все мальчики являются правонарушителями, очевидно, было бы неверно. Отсюда коэффициент Φ , формула которого приспособлена для того, чтобы отразить эту взаимную зависимость, равен всего лишь 0,5. Связь является взаимной, но не идеальной (абсолютной).

Несомненно, коэффициенты Q и Φ измеряют различные

аспекты взаимосвязи в четырехклеточной таблице. По существу, различие между двумя показателями заключено в том факте, что формула для Φ позволяет отразить степень двусторонней взаимосвязи с помощью единичного индекса. Коэффициент Φ «схватывает» любую двустороннюю взаимосвязь, которая существует между двумя рядами значений качественных признаков. Так как данный показатель измеряет только двустороннее отношение между x и y , то он в равной мере учитывает воздействие обеих переменных, поэтому его можно назвать обратимым.

Это можно очень просто проиллюстрировать (см. Табл. 4.5) в случае совершенной двусторонней взаимосвязи – ситуации, в которой нет отклонений ни в ту, ни в другую сторону: все мальчики являются правонарушителями, а все правонарушители – мальчиками.

Таблица 4.5

Пример идеальной двусторонней связи между полом и правонарушениями

	Мальчики	Девочки	Итого
Правонарушители	20	0	20
Не-правонарушители	0	80	80
Итого	20	80	100

В данном случае коэффициент контингенции также равнялся бы единице:

$$\Phi = \frac{1600 - 0}{\sqrt{20 \cdot 80 \cdot 20 \cdot 80}} = \frac{1600}{1600} = 1.$$

Вследствие этой совершенной двусторонней взаимосвязи, наблюдения распределяются вдоль одной диагонали матрицы 2×2 , а клетки, расположенные на другой диагонали, остаются пустыми. Поскольку каждый признак всецело «объясняет» другой, то Φ должен быть равен единице. Если выразить это в арифметической форме, то два нуля на одной из диагоналей таблицы с необходимостью приводят к $\Phi = 1$. Этот принцип обратимости также хорошо выдерживается для любого промежуточного

значения от нуля до единицы, то есть когда спаренные признаки только частично объясняют друг друга. По мере того как степень взаимной зависимости снижается, значение Φ также уменьшается.

Коэффициенты Q и Φ , конечно, совпадают при наличии идеальной двусторонней взаимосвязи по той простой причине, что полная односторонняя связь ($Q = 1$) является необходимым элементом совершенной двусторонней связи ($\Phi = 1$).

4.6. Коэффициенты связи, основанные на моделях прогноза (для номинальных признаков)

Что имеется в виду, когда говорится о прогнозных математических моделях? Ответ прост: в данном случае речь идет о взаимосвязанных (коррелирующих) признаках, однако связь эта такова, что наблюдение за «поведением» одного признака дает возможность предвидеть «поведение» другого. Или другими словами: признаки будут считаться взаимосвязанными, если реализованное значение одного из них позволяет относительно точно предугадать, каким будет значение другого [3, с. 143–180; 7, с. 327].

Например, если мы хотим опросить людей предпенсионного возраста на предмет их отношения к новой пенсионной реформе, нам и в голову не придет искать респондентов в ночном клубе или диско-баре, так как мы точно знаем, что средний возраст посетителей подобных увеселительных заведений не превышает тридцати лет. Получается, что признак «1» – «место нахождения» находится в тесной взаимосвязи с признаком «2» – «возраст». Более того, не зная значений признака «2» (конкретных возрастных показателей), мы можем с большой степенью вероятности их угадать, зная значения признака «1» (конкретные заведения).

Однако в данном случае правомерность сделанного нами вывода подтверждается, скорее, обыденным знанием, чем научным. Тем не менее есть вполне научные методы, с помощью которых, используя соответствующие формулы и осуществив

необходимые вычисления, можно спрогнозировать ситуацию, обосновав этот прогноз с помощью конкретных (относительно точных) цифр и показателей.

Для измерения такого рода корреляционной зависимости могут быть применены меры λ (лямбда) Гуттмана. Следует отметить, что в специальной литературе можно встретить формулы вычисления коэффициента λ Гуттмана, записанные как через абсолютные частоты, так и через частоты. В данном учебнике мы решили не приводить их все, чтобы не утяжелять текст излишними формулами, что, на наш взгляд, усложнит восприятие и усвоение самой важной информации.

Кроме того, важно помнить, что фамилия Гуттмана в разных учебниках, учебных пособиях и другой специальной научной литературе приводится по-разному, например: *Гуттмен*, *Гуттман*, *Гудман*. Такие расхождения связаны с «трудностями» перевода. Надо понимать, что если речь идет о коэффициенте корреляции и при этом упоминается какой-либо вариант фамилии из списка, представленного выше, – на самом деле имеется в виду один и тот же человек, известный всему миру социолог, статист, психолог Луи Гуттман¹.

¹ **Louis (Eliyahu) Guttman** (Бруклин, 1916 – Минеаполис, 1987) известный ученый, прославившийся своей изощренностью в применении математических методов в социологии и психологии. Блестящий новатор, который «видел теорию в методе и метод теории». Один из самых влиятельных представителей психометрического анализа 20-го столетия. В 1942 получил Степень Ph.D. (доктора философии) по социологии по теме, связанной с факторным анализом (Университет Миннесоты). В 1947 иммигрировал в Израиль, где основал и возглавил израильский Институт прикладного социального исследования (позже – Институт Гуттмана). С 1955 – профессор Социальной и Психологической Оценки в Еврейском Университете Иерусалима, где работал практически до последних дней жизни. Авторству Гуттмана принадлежит множество публикаций и научных книг по социологии, психологии и статистике, многие из которых получили всемирную известность.

Вообще, существует три меры λ Гуттмана: две из них – направленные, а одна представляет собой усреднение первых двух. Направленная мера показывает силу зависимости поведения одного признака от поведения другого. Например, направленная мера Гуттмана λ_{yx} («лямбда_игрик_по_икс») показывает, как изменяются значения признака Y в зависимости от изменения значений признака X , то есть, по сути, силу влияния X на Y . Формула направленной меры (коэффициента) Гуттмана λ_{yx} выглядит так:

$$\lambda_{yx} = \frac{\sum \max n_{ij} - \max n_{ci}}{n - \max n_{ci}},$$

где • $\max n_{ij}$ – максимальные частоты по строкам;
 • $\max n_{ci}$ – максимальная маргинальная частота по столбцам;
 • n – объем выборочной совокупности.

Алгоритм вычисления второй направленной меры Гуттмана λ_{xy} («лямбда_икс_по_игрик») аналогичен приведенному выше, просто меняются местами столбцы и строки таблицы сопряженности, а сам коэффициент показывает, как изменяются значения признака X , в зависимости от изменения значений признака Y .

Третья, усредненная, мера показывает общую силу связи между признаками и вычисляется довольно просто – путем нахождения среднего арифметического из двух направленных мер. Соответствующая формула вычисления коэффициента λ Гуттмана имеет следующий вид:

$$\lambda = \frac{\lambda_{yx} + \lambda_{xy}}{2}.$$

Теперь, когда формулы известны, рассмотрим вычисление всех трех мер (коэффициентов) λ Гуттмана на конкретном примере и попытаемся сделать научно обоснованный прогноз.

Представим себе гипотетическую ситуацию измерения связи между профессией человека и уровнем удовлетворенности жизнью. Попытаемся определить, сможем ли мы сделать прогноз о том, насколько человек удовлетворен своей жизнью, зная только его профессию (λ_{xy}). Или наоборот, сможем ли мы отгадать

профессию человека, зная только то, насколько он удовлетворен своей жизнью. В таблице 4.6 представлены абсолютные частоты двумерного распределения по обозначенным признакам. Для облегчения вычислительных процедур в эту таблицу мы добавили дополнительный столбец (с максимальными частотами по строкам) и дополнительную строку (с максимальными частотами по столбцам).

Таблица 4.6

**Результаты двумерного распределения по признакам
«Профессия» и «Удовлетворенность жизнью»**

Профессия (X)	Количество респондентов по уровню удовлетворенности жизнью (Y)			Максимальная частота	Общее число респондентов
	1	2	3		
Учитель	5	8	18	18	31
Продавец	12	7	14	14	33
Повар	7	11	20	20	38
Водитель	19	10	6	19	35
Максимальная частота	19	11	20		50
Итого	43	36	58	71	137 (n)

Подставив данные этой таблицы в соответствующие формулы получим:

$$\lambda_{yx} = \frac{(18 + 14 + 20 + 19) - \max 43; 36; 58}{137 - \max 43; 36; 58} = \frac{71 - 58}{137 - 58} \approx 0,16.$$

$$\lambda_{xy} = \frac{(19 + 11 + 20) - \max 31; 33; 38; 35}{137 - \max 31; 33; 38; 35} = \frac{50 - 38}{137 - 38} \approx 0,12.$$

$$\lambda \approx \frac{0,16 + 0,12}{2} \approx \frac{0,28}{2} \approx 0,14.$$

Результаты вычислений показывают, что в целом связь между признаками довольно слабая. Это касается всех трех мер λ Гуттмана. Даже по тому, как вычисляется коэффициент, видно, что он позволяет определять, существуют ли в строках модальные

группы, то есть есть ли в каждой профессиональной группе ярко выраженная, часто встречаемая «степень удовлетворенности жизнью». Судя по нашей таблице, таких групп практически нет, что и подтверждается маленьким значением коэффициента.

Какими же свойствами обладает этот коэффициент? Его значения могут варьироваться от «0» до «1», соответствующее неравенство имеет следующий вид: $0 \leq \lambda \leq 1$.

В случае, когда значение коэффициента равно «1», вероятность статистического предсказания Y по X максимальная. Такой случай практически в социологических исследованиях не встречается. Если рассмотреть его на нашем примере, то равным единице коэффициент λ Гуттмана мог бы быть только в том случае, если бы в каждой профессиональной группе все респонденты имели одинаковую степень удовлетворенности жизнью, и при этом в каждой из групп эта степень удовлетворенности была бы «своей».

Значение λ Гуттмана, равное нулю, свидетельствует о том, что знание признака X нечего не даст нам для увеличения знания о признаке Y и/или наоборот.

Представляется важным отметить, что в реальных исследованиях значения коэффициента Гуттмана очень малы и использовать их нужно, как и многие другие коэффициенты, в сравнительном контексте. Например, для ранжирования «как бы» независимых между собой признаков по степени их влияния на какой-то конкретный, особенно важный для исследователя, признак (обозначаемый как целевой или зависимый).

4.7. Корреляция рангов

Измерение с помощью рангов. Переменные могут быть связанными таким образом, что вариация одной переменной соответствует вариации другой и тогда о них говорят, что они ковариантны. Уровень рождаемости и уровень дохода в семье, число самоубийц и доля верующих – все это может быть связано друг с другом, а степень связи измерена с помощью подходящего

показателя корреляции. Из всего множества известных корреляционных индексов мы обратимся к рассмотрению только двух – коэффициента парной ранговой корреляции Спирмена (r_s) и коэффициента конкордации (W), который определяет степень множественной связи между качественными признаками.

Коэффициент ранговой корреляции используется для измерения взаимозависимостей между качественными признаками, значения которых могут быть упорядочены или проранжированы по степени убывания (или нарастания) данного качества у исследуемых социальных объектов.

Самый простой способ упорядочения данных состоит в их ранжировании. Простота этого способа заключается в том, что измерение может производиться интуитивно и субъективно, а не посредством явных, объективных единиц измерения. Такое упорядочение не может быть очень точным, и все же оно иногда оказывается очень полезным. В сущности, многие социологические понятия и не могут быть упорядочены никаким иным образом. Значительное число социологических исследований базируется на подобной субъективной основе. Так, профессии могут быть ранжированы по престижу; коллектив студентов может быть ранжирован в порядке предпочтения; национальности могут быть расположены в ряд посредством благосклонного или неблагосклонного к ним отношения; художественные картины могут быть ранжированы по эстетической ценности. Даже когда единицы измерения существуют как, например, при определении размеров городов или величины рождаемости, к ранжированию можно прибегнуть в том случае, если подобная точность не требуется. Так как ранговый порядок основывается на принципе «больше или меньше», можно рассматривать совокупность упорядоченных наблюдений как количественную, а не качественную переменную. Когда два ряда рангов изменяются совместно, то можно говорить о ранговой корреляции.

Коэффициент ранговой корреляции Спирмена (r_s). Простейший случай корреляции рангов – корреляция только между двумя рядами рангов. Когда число ранжированных объектов не превышает шести, наблюдатель в состоянии путем

простого обозрения вывести грубое, но приемлемое заключение о степени связи между ними. Рассмотрим два ряда расположения шести картин двумя экспертами, как показано в таблице 4.7.

Таблица 4.7

Пример полной положительной корреляции

Картина	Эксперт X	Эксперт Y	Различия (d)
A	1.	1.	0
B	2.	2.	0
C	3.	3.	0
D	4.	4.	0
E	5.	5.	0
F	6.	6.	0

Поскольку эти ряды дублируют друг друга, то различие между каждой парой рангов равно нулю. Кроме того, имеется полное соответствие одного рангового порядка другому. Таким образом, мерой корреляции между рангами должна быть единица. С другой стороны, ряд рангов можно было бы упорядочить от низшего уровня к высшему, никак не повлияв на степень предсказуемости, вместо того, чтобы сделать это упорядочение от высшего уровня к низшему (см. Табл. 4.8).

Таблица 4.8

Пример полной отрицательной корреляции

Картина	Эксперт X	Эксперт Y	Различия (d)	d^2
A	1.	6	-5	25
B	2.	5.	-3	9
C	3.	4.	1	1
D	4.	3.	1	1
E	5.	2.	3	9
F	6.	1.	5	25
Итого			0	70

Несмотря на то, что два ряда рангов находятся сейчас по отношению друг к другу в строго обратном порядке, отклонение каждого ранга от среднего ранга остается неизменным, и поэтому взаимная предсказуемость их значений остается той же самой.

Отсюда мерой связи все еще является единица, но уже с отрицательным знаком.

Третий возможный вид зависимости – между этими двумя крайностями – случайная связь, при которой ранг X нельзя предсказать по рангу Y (см. Табл. 4.9); любой X -ранг с равной степенью вероятности мог бы сочетаться с любым Y -рангом. В этом случае, по крайней мере, теоретически, корреляция была бы равна нулю.

Таблица 4.9

Пример случайной связи

Картина	Эксперт X	Эксперт Y	Различия (d)	d^2
A	1.	3	–2	4
B	2.	5.	–3	9
C	3.	1.	2	4
D	4.	6.	–2	4
E	5.	2.	3	9
F	6.	4.	2	4
Итого			0	34

Формула для измерения степени корреляции между двумя рядами рангов, приспособленная к условному диапазону колеблемости значений коэффициента от 0 до 1, была выведена Спирменом в 1904 году и имеет следующий вид:

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

где d – разность между парными рангами,

N – число ранжированных объектов.

Решая три вышеприведенных примера, получим следующие коэффициенты:

Таблица 4.7

$$r_s = 1 - \frac{6 \cdot 0}{6(36 - 1)} = 1$$

Таблица 4.8

$$r_s = 1 - \frac{6 \cdot 70}{6(35)} = -1$$

Таблица 4.9

$$r_s = 1 - \frac{6 \cdot 34}{6(35)} = 0,03$$

Как видно, величина может изменяться в пределах от -1 до $+1$, когда два ряда проранжированы в одном порядке.

Соответствующее неравенство имеет следующий вид: $-1 \leq r_s \leq 1$. При полном взаимном беспорядочном расположении рангов $r_s = 0$.

Значимость коэффициента корреляции Спирмена можно определить по таблицам критических величин r_s (см. Приложение, табл. «А», «Г»). Наблюдаемые значения критерия вычисляются по следующей формуле:

$$z = \frac{r_s}{\sqrt{N-1}}.$$

Следует отметить, что социологи наиболее часто прибегают к вычислению коэффициента корреляции Спирмена в тех случаях, когда необходимо сравнить две различные выборки. Например, когда мы хотим сравнить, насколько отличаются/совпадают ценностные ориентации студентов первых курсов в двух различных вузах или жизненные планы старшеклассников нынешнего и прошлого года (однако при условии, что соответствующие измерения по обозначенным признакам на разных выборках осуществлялись с помощью одних и тех же шкал).

Техника вычисления в случае объединенных рангов. Ранги иногда могут объединяться. Эксперт в затруднительном положении может оценить две картины одинаково, количественные значения и меры могут также быть одинаковыми. В подобных случаях два или более наблюдений могут, по видимому, одновременно претендовать на один и тот же ранг. Так как число рангов и число наблюдений должно совпадать, то просто невозможно двум наблюдениям присвоить один и тот же ранг, поэтому им должны быть присвоены объединенные ранги. В таких случаях обоим наблюдениям приписывается значение среднего арифметического из двух объединенных рангов. Так, если объединяются 3-е и 4-е наблюдения, каждому присваивается ранг 3,5. Если объединяются три (и более) ранга, действует то же самое правило усреднения. Поскольку такие объединения затушевывают различия рангов, они приводят к уменьшению предельного значения корреляции.

Корреляция между упорядоченными переменными. Если даны относительно короткие ряды количественных данных

и если нет необходимости в свойственной им точности, можно расположить эти данные в порядке их величин и коррелировать ранги вместо числовых значений. Преимущества этого приема заключаются в скорости и простоте вычисления, что компенсирует уменьшение точности корреляционного показателя. В таблице 4.10 коррелируются доход и месячная квартплата шести семей посредством сравнения ранговых порядков.

Таблица 4.10

Ранжирование семей по доходу и месячной квартплате

<i>Семья</i>	<i>Доход</i>	<i>Кв. плата</i>	<i>Ранг</i>
А	1500	250	1
Б	700	200	2
В	500	125	3
Г	450	90	4
Д	400	80	5
Е	300	75	6

Как видно из таблицы, столбцы значений находятся в полном ранговом соответствии: ранговая градация по доходу точно соответствует градации по квартплате. Однако исходные величины не обнаруживают такого же полного соответствия. Конечно, такие ряды спаренных рангов дадут r_s , равное единице, но только потому, что игнорируют детали исходной информации. Практически, для доказательства того факта, что корреляция между количественными данными не равна 1, будет вполне достаточно для построения корреляционного поля, поскольку некоторые точки на нем будут отклоняться от прямой линии, которая изображает полную линейную корреляцию.

Анализ $r_s \times r_s$ можно действительно рассматривать в качестве меры конгруэнтности (или согласия), колеблющейся от полного совпадения (+1) через случайную связь (0) к полному несоответствию (–1) между рядами рангов. Отсюда очевидна обоснованность базирования формулы на различиях между спаренными рангами.

Однако различия между рангами имеют более точный статистический смысл, который не виден из формулы. Но студент

должен уже сознавать, что формулы часто являются в высшей степени конденсированными, операциональными конструкциями, которые скрывают больше, чем обнаруживают. Поэтому совсем не очевидно, что r_s коррелирует скорее с сигма-значениями (σ), а не с рангами. Уже имели возможность заметить, что ранговая корреляция основывается не на прямом соответствии между ранговыми порядками, а скорее на соответствии между отклонениями соответствующих рангов от среднего ранга. Эти отклонения измеряются в сигма-единицах. Формула автоматически превращает ранги в σ и отсюда действительные различия между рангами (d) – в различия, выраженные в сигмах (σ). Следовательно, можем подставлять в нее исходные данные, упорядоченные по рангам.

Однако вышеупомянутая серия операций основывается на предположении, что интервалы между рангами равны. Например, предполагается, что разрыв между рангами 1 и 2 равен разрыву, который существует между рангами 2 и 3. И все же здравый смысл подсказывает, что идеального случая равных интервалов, которые требуются для формулы, достигнуть невозможно.

Действительная степень субъективного предпочтения первого ранга по сравнению со вторым не является обязательно такой же по интенсивности, как предпочтение второго ранга по сравнению с третьим. Так, лошадей можно расположить в порядке пересечения ими финишной линии, но интервалы между ними не будут одинаковыми. Тем не менее, r_s может использоваться и используется в том случае, когда отсутствуют объективные единицы измерения или когда различие между неравными интервалами рассматривается как несущественное.

Из вышесказанного становится понятным, что r_s измеряет корреляцию между порядковыми рангами, а не между ранжируемыми величинами. Следовательно, r_s преувеличивает степень связи между изучаемыми переменными. Таким образом, два эксперта, упорядочивающие одну и ту же совокупность художественных полотен, могут иметь различный художественный вкус и все же расположить произведения искусства в идентичном порядке. Эксперт 1 может считать все их, в высшей степени,

превосходными, тогда как эксперт 2 будет оценивать одинаково низко всю совокупность картин. Эксперты могут согласиться относительно порядка, но разойтись во мнении по существу. Несмотря на все это, для большинства реальных ситуаций разумно предположить, что ввиду общности культуры сходное ранжирование всегда соответствует сходным предпочтениям. Когда два ряда данных не располагаются на единой непрерывной шкале, как в случае арендной платы и дохода, или когда они относятся к несоизмеримым категориям, как, например, рождаемость и доход, то подобная проблема обоснования их расположения на шкале, естественно, не возникает. И все же, прежде чем использовать r_s , исследователь должен убедиться в том, что его интересует корреляция только между рангами, а не между фактическими величинами.

Мера соответствия для трех и более ранговых рядов.

Формула для вычисления r_s пригодна только для двух ранговых рядов, однако данные могут состоять из трех и более ранговых рядов. Один из методов определения общей степени соответствия между тремя и более упорядоченными рядами данных состоит просто в вычислении среднего арифметического из всех возможных значений r_s . Таким образом, если бы три эксперта ранжировали шесть картин, можно было бы вычислить r_s для всех возможных сочетаний рядов по парам для того, чтобы определить среднее согласие между ними. Спаренные ряды сочетались бы следующим образом: эксперты 1 и 2, 2 и 3, 1 и 3. Затем три значения усредняются при соблюдении знаков. Результат подобной операции называется иногда сводной взаимокорреляцией рангов; однако его можно было бы более строго назвать по причинам, указанным ниже, коэффициентом соответствия (конкордации). Коэффициент конкордации используется для измерения степени согласованности двух или нескольких рядов проранжированных значений переменных.

В таблице 4.11 представлен результат ранжирования тремя экспертами каких-то условных предметов. Средняя величина из трех значений r_s показывает лишь умеренную степень соответствия. Теперь предположим, что имеются три значения r_s : 1; -1 и -1.

Средняя величина из трех полных корреляций равна вовсе не 1,00, как можно было бы наивно полагать, а всего лишь – 0,33. Все это становится понятным только тогда, если рассматривать эту среднюю как выражение степени соответствия, а не корреляции. В данном случае, при трех полных корреляциях, преобладающим отношением является отношение несоответствия при одном полном соответствии.

Таблица 4.11

**Результат ранжирования экспертами
неких «условных параметров»**

<i>Эксперты</i>	<i>1 и 2</i>	$r_s = 1$
<i>Эксперты</i>	<i>2 и 3</i>	$r_s = 1$
<i>Эксперты</i>	<i>1 и 3</i>	$r_s = -1$

Для табличных данных $\bar{r}_s = -0,33$.

Хотя усреднение значений показателей любого типа обычно сопряжено с возможными ошибками, в данном случае оно оправдано по той причине, что все значения r_s имеют равный вес, то есть все они вычислены на одном и том же числе наблюдений. Следует повторить, однако, что сама эта средняя величина не является коэффициентом корреляции; она есть средняя из нескольких коэффициентов. Эта величина никогда не может принять значение «минус единица» по той простой причине, что если два ряда коррелируют между собой со значением -1 , то третий не может коррелировать с ними обоими одинаковым образом. Однако средняя величина может стать равной «плюс единица» в случае полного соответствия между всеми рядами.

Когда усредняется очень много значений r_s , то вычисления становятся все более трудоемкими. Тем не менее разработан очень простой метод усреднения значения r_s ; этот метод особенно полезен, когда число сравниваемых рядов велико. Формула эта не так громоздка, как кажется на первый взгляд, поскольку все символы в ней являются общепринятыми величинами.

$$W = \frac{12S}{k^2 \cdot N \cdot (N^2 - 1)},$$

где k – число переменных;

N – число индивидов или категорий, которые ранжируются;

$S = (\text{Сумма рангов по строке } \underline{\text{минус}} a)^2$, a – среднее из суммы рангов.

В приведенной ниже таблице 4.12 демонстрируется применение этой формулы.

Таблица 4.12

Пример вычисления множественного коэффициента ранговой корреляции (коэффициента конкордации W)

Респондент	Удовлетворенность по признакам А, Б, В			Сумма рангов
	А	Б	В	
1-й	1	2	1	4
2-й	3	4	5	12
3-й	5	5	4	14
4-й	4	3	3	10
5-й	2	1	2	5
$N = 5$				$\Sigma = 45$

Для данных таблицы $a = 45/5 = 9$;

$$S = (4 - 9)^2 + (12 - 9)^2 + (14 - 9)^2 + (10 - 9)^2 + (5 - 9)^2 = 76;$$

$$W = \frac{12 \times 76}{3^2 \times 5 \times (5^2 - 1)} = 0,84.$$

Значимость полученной величины W для $N > 7$ проверяется

по критерию χ^2 : $\chi^2 = \frac{12S}{kN(N-1)}$ со степенью свободы $N - 1$.

Для рассматриваемого примера $\chi^2_{набл} = 10,133$, степень свободы $(N - 1) = 4$. Для $\alpha = 0,05$ по таблице критических значений χ^2 находим $\chi^2_{кр} = 9,488$ (см. Приложение 2, табл. Б). Поскольку наблюдаемое значение χ^2 больше критического, то необходимо отвергнуть гипотезу H_0 о том, что не существует значимой связи между рассматриваемыми переменными.

Полезность r_s . Так как r_s применяется к порядковым данным,

не имеющим определенных единиц измерения, этот показатель весьма полезен для социолога, которому часто приходится иметь дело с данными, носящими субъективный характер. Социометрическое ранжирование предпочтений, эстетических суждений и других аналогичных явлений – все это можно коррелировать для того, чтобы определить степень согласия в предпочтениях выборщиков и оценках экспертов. Профессии могут быть ранжированы как по социально-экономическому уровню, так и по уровню разводов. Следовательно, тем самым может быть получена мера предполагаемой взаимосвязи между социально-экономическим уровнем и частотой разводов.

4.8. Линейная корреляция (для метрических признаков и интервального уровня измерения)

Необходимость общей меры корреляции. Корреляционное поле, которое было тщательно рассмотрено в предыдущем подразделе, позволяет визуально обнаружить совместные вариации двух переменных. Группы изображенных на корреляционном поле точек, которые концентрируются вокруг какой-то гипотетической линии, наводят на мысль об определенном «законе связи» между двумя переменными. Более того, наблюдая ширину разброса, можно сделать, по крайней мере, предварительное заключение о том, насколько хорошо события соответствуют этому гипотетическому закону. Такие заключения зачастую имеют большое значение, но все же оставляют желать лучшего, поскольку субъективны и нестандартизованы. Следовательно, их нельзя как-либо описать или связать с чем бы то ни было, не воспроизводя корреляционного поля. А поскольку эти «визуальные критерии» не точны в математическом смысле, их сопоставление невозможно, даже если в распоряжении исследователя имеются все корреляционные поля во всей их сложности. Поэтому необходим объективный, стандартный, синоптический критерий связи между двумя переменными, конкретная мера корреляции.

То изображение, которое мы видим на корреляционном поле, является гипотетической линией тенденции, которая более или менее отображает совокупность точек, и на основании которой оценивается степень корреляции. Но для того чтобы измерить эту корреляцию, необходимо:

- а) точно установить положение такой линии;
- б) измерить степень согласия событий, которые ее составляют;
- в) посредством специального метода перевести этот результат в индекс корреляции.

Итак, прежде всего, следует позаботиться об определении расположения линии, которая наилучшим образом аппроксимирует наблюдаемые данные.

Напомним, что по типу корреляционная связь может быть прямой или обратной. Прямая связь означает, что при увеличении значения одного признака в среднем увеличивается значение другого, а обратная – при увеличении одного признака в среднем уменьшается значение другого.

По форме корреляционная связь может быть прямолинейной или криволинейной. Гипотезу о форме связи устанавливают по корреляционному полю. **Прямолинейную** форму связи имеет такая связь, при которой с увеличением фактора результирующий признак увеличивается или уменьшается на одну и ту же величину. **Криволинейная** связь наблюдается тогда, когда подобное изменение происходит неравномерно.

Если связь прямолинейна, то ее можно выразить с помощью уравнения прямой $y = kx + b$, если же связь криволинейна, то она может быть выражена посредством любой кривой, как элементарной, так и нет, например, параболы ($y = ax^2 + bx + c$) или гиперболы ($y = \frac{a}{x} + b$).

По тесноте корреляция может быть *тесной* или *слабой*. Под теснотой связи понимается мера, которая показывает, насколько чувствителен результирующий признак к изменениям факторного признака. О тесноте связи можно судить по корреляционному полю. Если точки распределения плотно сконцентрированы вдоль какой-либо кривой, то говорят, что связь **тесная**,

в противном случае делается вывод о *слабой* связи между переменными.

Корреляция может быть также *парной* или *множественной*. **Парная** связь устанавливается между двумя признаками (факторным и результирующим). **Множественная** связь устанавливается между большим количеством факторных признаков и результирующим.

Все характеристики корреляционного анализа, применяемого для объектов, измеренных по интервальной или порядковой шкале для количественных признаков, определяются следующими коэффициентами:

- 1) коэффициент наклона линии регрессии $R_{y/x}$;
- 2) коэффициент детерминации $-r^2$ и, соответственно, коэффициент недетерминированности $(1 - r^2)$;
- 3) коэффициент корреляции Пирсона $-r$;
- 4) корреляционное отношение η^2 .

Уже говорилось ранее, что связь может быть прямолинейной или криволинейной, следовательно, линия наилучшего приближения может быть прямой или кривой. В данном подразделе ограничимся рассмотрением только линейных связей, то есть таких, которые могут быть представлены в виде прямой линии.

После того как построено корреляционное поле, можно начертить прямую линию, которая при ближайшем рассмотрении оказалась бы линией основной тенденции. Если бы была проведена такая линия, то ее расположение можно было бы обосновать тем, что проведена она, как раз, через середину самой «густой» части скопления всех точек (значений), то есть по возможности, наиболее близко ко всем точкам в среднем, разрезая совокупность этих точек на две одинаковые по длине и ширине полосы. Линия была прочерчена вдоль середины, потому что интуитивно чувствовалось, что именно таким образом ее можно более всего приблизить ко всем наблюдаемым точкам в среднем. Таким образом, она служит оценкой для значений y_i переменной Y и учитывает каждое наблюдаемое значение x_i переменной X . Оценочные значения y_i соответствуют наблюдаемым значениям x_i и составляют вышеупомянутую линию тенденции. Естественно, чем

ближе расположение наблюдаемых точек к оценочной линии, то есть – чем меньше они отклоняются от нее, тем меньше ошибка предсказания y_i по отношению к x_i . Если, например, доход был бы единственным детерминантом социального статуса, то наблюдаемые точки, все без исключения, обязательно расположились бы на линии соответствующего корреляционного поля; если бы только лишь социальный статус определял доход, то все точки снова попали бы на данную линию.

Итак, какое же значение для корреляции имеют эти отклонения от гипотетической линии регрессии? Они указывают социологу на то, что помимо дохода существуют другие факторы, определяющие социальный статус. Действие этих других, неизвестных, факторов искажает предсказание о природе явления. Как следствие, корреляционная связь между двумя этими переменными оказывается слабой.

Отклонения от среднего арифметического. Сравнивая значения переменных при корреляции, исследователь должен привыкнуть мыслить в терминах отклонений от средних арифметических, вместо того, чтобы анализировать необработанные данные. Подобно многим другим статистическим операциям, которые могли бы, на первый взгляд, показаться излишне косвенными и сложными, эта мера также согласуется со здравым смыслом. Обычно не сравнивают сырые, необработанные результаты наблюдений (рождаемость, жалование и т. д.), а предпочитают оценивать их как «высокие» или «низкие», то есть как такие, которые лежат выше или ниже среднего показателя или привычной нормы. Однако нельзя сравнивать единицы различных категорий, таких как доход и рождаемость, в их исходной форме, поэтому связывают доход ниже среднего с размером семьи, который превышает среднее.

Таким образом, видно, что среднее арифметическое является удобной точкой начала отсчета, естественным стандартом, в случае, когда сравнивают два ряда данных в отношении их близости друг к другу.

Объясняемые и необъясняемые остаточные отклонения. Проведенная от руки строго через середину линия тенденции,

очевидно, является средней для всех наблюдаемых точек, непрерывной последовательностью средних значений, как бы приспособленной к всевозможным значениям случайных событий. Поскольку она состоит из ожидаемых отклонений, значения на наклонной линии тенденции в различных случаях называют «ожидаемыми», «предсказуемыми» или «оценочными» величинами. Так, для примера, рассматривающего взаимосвязь количества высших учебных заведений III–IV уровня аккредитации Украины и численности студентов в них, можно установить на языке статистики, что те отклонения, которые могут быть больше или меньше, чем наблюдаемые, объясняются количеством вузов или могут быть приписаны действию фактора количества вузов. Поэтому данные отклонения получили название объясняемых отклонений.

Поскольку имеются другие факторы, которые влияют на численность студентов (такие, как уровень рождаемости, оплата за обучение и т. д.), то наблюдаемые (действительные, эмпирические) величины не совпадают с ожидаемыми (теоретическими). Воздействие этих неизвестных факторов измеряется расхождениями между наблюдаемыми и ожидаемыми отклонениями. Чем меньше эти расхождения, тем слабее должно быть воздействие посторонних факторов, чем больше эти расхождения, тем сильнее должно быть это воздействие. Расхождения такого рода иногда называют остатками, которые необъяснимы на основании используемых данных, то есть они должны быть приписаны не количеству вузов, а неизвестным факторам, о сущности которых можно лишь делать предположения. Поэтому их называют необъяснимыми остаточными отклонениями; они также измеряются относительными расстояниями от линии тенденции до отдельных точек исследуемой совокупности. Чем ближе соответствие между объясняемыми и наблюдаемыми отклонениями, то есть чем меньше остатки, тем выше степень корреляции X и Y . Следовательно, было бы логично измерять степень корреляции в зависимости от степени соответствия между наблюдаемыми и ожидаемыми отклонениями или в виде той доли полного наблюдаемого отклонения, которая объяснима.

В этом и заключается основной принцип измерения корреляции, который всегда подсознательно используется при ее вычислении. Конечно, измерение можно было бы произвести с помощью линии регрессии, проведенной от руки, или путем графического анализа при помощи циркуля и линейки, однако существенным остается сравнение, или сопоставление, *полного* и *объясняемого* отклонений, а также нахождение необъясняемого остатка.

Поскольку эти понятия играют такую большую роль в последующем изложении, резюмируем их смысл более точно: 1) ***полное*** (общее или суммарное) наблюдаемое отклонение – это отклонение наблюдаемой величины от среднего арифметического значения; 2) ***объясняемое*** отклонение – это отклонение ожидаемого (регрессионного) значения от среднего арифметического значения; 3) ***необъясняемый остаток*** – это расхождение между полным и объясняемым отклонениями. Разумеется, это есть разность между вышеупомянутыми полным и объясняемым отклонениями.

Измерение линейной корреляции. Серьезный недостаток графиков, выполненных от руки, состоит в том, что они зависят от индивидуального суждения. Без использования стандартной техники вычислений маловероятно, чтобы два исследователя когда-нибудь расположили линию тенденции одинаковым образом. Очевидно, линия, которая используется для измерения корреляции предложенным выше способом, должна всегда удовлетворять одним и тем же требованиям, в противном случае результатам будет не хватать надежности, которая существенна в научной работе.

Одно из таких требований было сформулировано следующим образом: линия располагается так, чтобы сделать равной нулю сумму вертикальных отклонений от этой линии, то есть сбалансировать суммы положительных и отрицательных расхождений. Таким образом, эта линия изображает рассеяние вокруг нее точно так же, как среднее арифметическое представляет совокупность в целом. И подобно среднему арифметическому, она минимизирует сумму квадратов отклонений, что соответствует принципу наименьших квадратов. Линию следует расположить так, чтобы

минимизировать сумму квадратов горизонтальных отклонений, или X -остатков. Однако так как эта линия дала бы ту же меру корреляции, что и линия, минимизирующая квадраты вертикальных отклонений, то необязательно чертить их обе. Для измерения линейной корреляции между двумя переменными достаточно и одной линии регрессии. Имея в виду эти свойства, ее называют линией наилучшего приближения; согласно критерию наименьших квадратов, она приближает рассеяние лучше, чем любая другая прямая. Проведение этой математической линии наилучшего приближения упрощается изображением событий как отклонений от средних арифметических. При соблюдении этих условий линия наименьших квадратов всегда будет проходить через начало отсчета, которое лежит на пересечении средних. Следовательно, ее легко вычертить, как только будет определен ее наклон, – тангенс угла наклона прямой по отношению к оси X .

Вычисление наклона линии регрессии. Вычисление наклона этой линии соответствует следующему уравнению, представленному здесь без объяснения:

$$R_{y/x} = \frac{\sum (x - \bar{x})(y - \bar{y})}{\sum (x - \bar{x})^2},$$

где • $R_{y/x}$ – наклон линии ожидаемых значений в зависимости от X (читается: « R_y по X »);

- $\bar{x}$ – среднее арифметическое по признаку X ;

- $\bar{y}$ – среднее арифметическое по признаку Y ;

- $\sum (x - \bar{x})(y - \bar{y})$ – сумма произведений парных отклонений значений переменных X и Y (читается: «сумма парных произведений»).

Для примера вычисления коэффициента, показывающего наклон линии регрессии, воспользуемся таблицей динамики сети высших учебных заведений, к которой мы неоднократно обращались в предыдущих подразделах, в частности, посвященных вычислению среднего квадратического отклонения (см. *Табл. 4.13*).

Это отношение задает наклон искомой линии наилучшего приближения и представляет собой среднее изменение значе-

Таблица 4.13

**Динамика сети высших учебных заведений Украины
III–IV уровня аккредитации и численности студентов в них**

Количество высших учебных заведений Украины III–IV уровня аккредитации (x_i)	Численность студентов в них (тыс.) (y_i)	$(x_i - \bar{x}) \cdot (y_i - \bar{y})$	$(x_i - \bar{x})^2$
158	718,8	–10568,4	7766,016
159	680,7	–7129	7590,766
232	645	–651,516	199,5156
255	617,7	–167,0719	78,76563
274	595	–108,016	777,0156
280	526,4	–2455,09	1147,516
298	503,7	–4937,2	2691,016
313	503,7	–6364,83	4472,266
Σ 1969	4791	–32047	24722,88
$\bar{X} = 246,125$	$\bar{Y} = 598,875$		

ния y_i при изменении значения x_i на единицу. Соответственно, чтобы найти среднее изменение в численности студентов на единичное изменение в количестве вузов, вычисляют парные произведения, возводят в квадрат отклонения по X и образуют отношение между их соответствующими суммами. Все эти действия отображены в таблице 4.13, представленной выше. Осуществив соответствующие вычисления, мы получили $R_{y/x} = -1,296$. Это и есть средняя скорость изменения численности студентов на единицу изменения количества вузов, то есть при каждом изменении количества вузов на единицу численность студентов уменьшается в среднем на 1,296 тыс. человек. Поскольку эта величина фиксирует положение линии регрессии относительно оси независимой переменной (в нашем случае – оси X), она называется наклоном. Как только наклон определен, можно вычертить линию наименьших квадратов и перейти к точному измерению корреляции.

Коэффициент детерминации. Существует несколько альтернативных определений коэффициента детерминации, однако в случае линейной регрессии все они эквивалентны. **Коэффициент детерминации** – это показатель того, насколько измене-

ния зависимого признака объясняются изменениями независимого. Если более точно – это доля дисперсии независимого признака, объясняемая влиянием зависимого. Вычисление этого коэффициента основывается на том очевидном принципе, что чем ближе объясняемые отклонения к общей вариации, тем больше доля объясняемой вариации и тем выше степень корреляции. Чем значительнее доля объясненной вариации, тем меньше роль прочих факторов.

Как статистическая операция, преобразование этих отклонений в отдельный индекс состоит в суммировании квадратов объясняемых отклонений и выражает вариацию как долю общей вариации, которую необходимо объяснить. Символически это можно представить так:

$$r^2 = \frac{\sum_i (\hat{y}_i - \bar{y})^2}{\sum_i (y_i - \bar{y})^2},$$

где • $\hat{y}_i$ – ожидаемое значение зависимой переменной, предсказанное по уравнению регрессии;

• y_i – наблюдаемое (эмпирическое) значение зависимой переменной;

• $\bar{y}$ – среднее значение зависимой переменной.

Коэффициент детерминации изменяется в диапазоне от 0 до 1. Соответствующее неравенство имеет следующий вид: $0 \leq r^2 \leq 1$.

При этом, если $r^2 = 0$ – это означает, что связь между переменными регрессионной модели отсутствует, и вместо нее для оценки значения выходной переменной можно с таким же успехом использовать простое среднее ее наблюдаемых значений. Когда же $r^2 = 1$ – имеет место соответствие идеальной модели, когда все точки наблюдений лежат точно на линии регрессии, то есть сумма квадратов их отклонений равна 0. На практике, если коэффициент детерминации близок к 1, это указывает на то, что модель работает очень хорошо (имеет высокую значимость), а если – к 0, то это означает низкую значимость модели, когда независимая переменная (X) плохо «объясняет»

поведение зависимой переменной (Y), то есть линейная зависимость между ними отсутствует. Очевидно, что такая модель будет иметь низкую эффективность.

Другими словами, r^2 представляет собою долю в общей вариации зависимой переменной Y , которая обусловлена изменением независимой переменной X , поэтому r^2 и называют коэффициентом детерминации. Так как объяснить можно не более, чем всю полную вариацию, то r^2 никогда не может превышать единицу, а практически всегда бывает меньше ее.

Полное вычисление r^2 показано в таблицах 4.14 и 4.15, где и объясняемые и наблюдаемые отклонения возводятся в квадрат и суммируются, что дает соответственно объясняемую и общую вариацию.

Таблица 4.14

Распределение взаимных частот, описывающих связь дохода в сотнях гривен и оценкой социального статуса

Доход в сотнях гривень (x)	Оценка социального статуса (y)	$(x_i - \bar{x})$	$(y_i - \bar{y})$	$(x_i - \bar{x})^2$	$(x_i - \bar{x})(y_i - \bar{y})$
3	3	-7	-3	49	21
7	5	-3	-1	9	3
11	7	1	1	1	1
14	6	4	0	16	0
15	9	5	3	25	15
$\bar{X} = 10$	$\bar{Y} = 6$	0	0	100	40

Исходя из приведенной выше формулы для определения наклона линии регрессии, в данном случае $R_{y/x} = \frac{40}{100} = 0,4$. Значит, уравнение теоретических распределений выглядит так $\hat{y} - \bar{y} = R_{y/x}(x_i - \bar{x})$.

Продолжение вычисления для таблицы 4.14

$\hat{y} - \bar{y}$	$(\hat{y} - \bar{y})^2$	$(y_i - \bar{y})$	$(y_i - \bar{y})^2$
$0,4 \times (-7) = -2,8$	7,84	-3	9
-1,2	1,44	-1	1
0,4	0,16	1	1
1,6	2,56	0	0
2,0	4,00	3	9
0	16	0	20

$$r^2 = 16 / 20 = 0,8$$

Таким образом, деля объясняемую вариацию 16 на общую вариацию 20, мы получили $r^2 = 0,8$. Это и есть окончательная мера степени связи между двумя переменными: она показывает, что на 80% вариация переменной Y объясняется линейной зависимостью от переменной X .

Обратимость r^2 . С одинаковым успехом могли бы вычислить как регрессию X от Y , так и регрессию Y по X и тем самым определить долю вариации X , объясняемую через Y . Действительно, такой результат был бы существенен для окончательной формулировки корреляции. Почему же в таком случае вычисляют r_{xy}^2 (X по Y)? Ответ состоит в том, что вклад r_{xy}^2 учитывается автоматически. Нет необходимости строить линии регрессии дважды по той причине, что $r_{xy}^2 = r_{yx}^2$. Короче говоря, r^2 обратимо.

Например, зная, что $r_{yx}^2 = 0,80$, можно утверждать не только, что доход объясняет 80% вариации в социальном статусе, но также и наоборот, а именно: социальный статус объясняет 80% вариации в доходе. Так как с точки зрения статистики одно объясняет другое в одинаковой степени, то подписные значки при r^2 обычно опускают.

Коэффициент недетерминированности. Поскольку разность между общей вариацией и объясняемой вариацией обязательно равна необъясняемой вариации, то необъясняемая

доля вариации есть просто разность от единицы и соответственно называется **коэффициентом недетерминированности**. Его можно записать следующим образом:

$$(1 - r^2) = 1 - \frac{\sum_i (\hat{y}_i - \bar{y})^2}{\sum_i (y_i - \bar{y})^2},$$

где • $\hat{y}_i$ – ожидаемое значение зависимой переменной, предсказанное по уравнению регрессии;

• y_i – наблюдаемое (эмпирическое) значение зависимой переменной;

• $\bar{y}$ – среднее значение зависимой переменной.

Можно бы получить эту величину непосредственным измерением остатков, расположенных вокруг линии регрессии, возведением их в квадрат и суммированием, выражая эту сумму как долю полной вариации. В действительности, в этом и состоит операциональный смысл коэффициента недетерминированности $1 - r^2$. Но независимо от того, вычисляется ли этот коэффициент непосредственно из остатков или косвенно через величину r^2 , его всегда можно истолковать как долю вариации зависимой переменной, которую нельзя линейно объяснить через независимую переменную. Поэтому он измеряет силу воздействия неизвестных факторов.

Коэффициент корреляции Пирсона (r). Важной характеристикой совместного распределения двух случайных величин является ковариация. В теории вероятности **ковариация** – это мера линейной зависимости случайных величин. Считается, что если две случайные величины X и Y имеют в отношении друг друга линейные функции регрессии, то эти величины (X и Y) ковариантны, а значит, связаны между собой линейной корреляционной зависимостью. Если двумерная случайная величина (X, Y) распределена нормально, то между X и Y имеет место линейная корреляция.

Если величины X и Y независимы, то $\text{cov}(X, Y) = 0$. Такие величины, то есть те, которые имеют нулевую ковариацию, назы-

ваются **некоррелированными**. Ковариация – мера связи, широко применяемая в физике и технических науках, которая определяется по следующей формуле:

$$COV_{xy} = \frac{\sum_i (x_i - \bar{x})(y_i - \bar{y})}{N - 1},$$

- где
- x_i – значения, принимаемые переменной X ;
 - y_i – значения, принимаемые переменной Y ;
 - $\bar{x}$ – среднее значение по X ;
 - $\bar{y}$ – среднее значение по Y ;
 - N – число единиц совокупности.

Однако в социологии, в отличие от физики, например, большинство переменных измеряется в произвольных шкалах. Ковариация служит характеристикой взаимозависимости случайных величин, однако только лишь по ее абсолютному значению нельзя судить о том, насколько сильно величины взаимосвязаны, так как величина ковариации зависит от единиц измерения независимых величин. Данная особенность ковариации затрудняет ее использование в целях корреляционного анализа. Для устранения недостатка ковариации, для того, чтобы сделать меру связи независимой от единиц измерения того или иного признака, достаточно разделить ковариацию на соответствующие стандартные отклонения [4, с. 70]. Таким образом и была получена формула **коэффициента корреляции Пирсона**, который разработали Карл Пирсон, Фрэнсис Эджуорт и Рафаэль Уэлдон в 90-х годах XIX века. Этот коэффициент рассчитывается по следующей формуле:

$$r = \frac{\sum_i (x_i - \bar{x})(y_i - \bar{y})}{(N - 1)\sigma_x \sigma_y}$$

или

$$r = \frac{\sum_i (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_i (x_i - \bar{x})^2 \times \sum_i (y_i - \bar{y})^2}},$$

где • x_i – значения, принимаемые переменной X ;
 • y_i – значения, принимаемые переменной Y ;
 • $\bar{x}$ – среднее значение по X ;
 • $\bar{y}$ – среднее значение по Y ;
 • σ_x, σ_y – выборочные средние квадратические отклонения величин X и Y .

Последняя формула является основной для вычисления коэффициента корреляции Пирсона. Этот коэффициент показывает тесноту линейной связи между X и Y : чем ближе $|r|$ к единице, тем сильнее линейная связь между X и Y . Следовательно, r изменяется в пределах от -1 до 1 , что может быть представлено следующим неравенством: $-1 \leq r \leq 1$.

В качестве примера вычисления вновь обратимся к данным таблицы 4.14 уже рассматриваемой ранее в этом подразделе.

Таблица 4.16

Вычисление для данных таблицы 4.14
«Распределение взаимных частот, описывающих связь дохода в сотнях гривен и оценкой социального статуса»

Доход в сотнях гривен (x)	Оценка социально-го статуса (x)	$(x_i - \bar{x})$	$(x_i - \bar{x})^2$	$(y_i - \bar{y})$	$(y_i - \bar{y})^2$	$(x_i - \bar{x})(y_i - \bar{y})$
3	3	3-10=-7	49	3-6=-3	9	(-7)×(-3)=21
7	5	-3	9	-1	1	3
11	7	1	1	1	1	1
14	6	4	16	0	0	0
15	9	5	25	3	9	15
$\bar{X} = 10$	$\bar{Y} = 6$		$\sum_i = 100$		$\sum_i = 20$	$\sum_i = 40$

Подставив зафиксированные в таблице 4.16 цифры в основную расчетную формулу для нахождения коэффициента корреляции Пирсона, получим:

$$r = \frac{40}{\sqrt{100 \times 20}} = \frac{40}{44,72} = 0,89.$$

Такое значение коэффициента свидетельствует о наличии весьма сильной прямой связи между уровнем дохода респондентов и их оценкой собственного социального статуса.

В буквальном смысле величину r можно истолковывать как среднее изменение y на каждое единичное изменение x , или наоборот, предполагая в обоих случаях, что критерии выражены в стандартной форме, то есть в долях σ . Итак, если известно, что доход отклоняется на величину 1σ от среднего, то следует ожидать, что соответствующий ему уровень социального статуса отклоняется в среднем от своего среднего арифметического значения на $0,89\sigma$.

На основании данного принципа можно оценить y (в стандартной форме) для любого данного значения x точно так же, как это было сделано для получения ожидаемых отклонений y .

Следует лишь применить формулу: $\hat{z}_y = r(z_x)$, где $\hat{z}_y$ – оценочное значение в единицах σ для величины x . Подставляя в формулу данные, получаем оценочные величины, приведенные в таблице 4.17.

Таблица 4.17

Оценка z_y по наблюдаемому z_x

Доход в сотнях гривень (x)	Оценка социального статуса (y)	z_x	$\hat{z}_y = r(z_x)$
3	3	-1,56	-1,39
7	5	-0,67	-0,60
11	7	0,22	0,20
14	6	0,89	0,79
15	9	1,12	1,00
$\bar{X} = 10$	$\bar{Y} = 6$		

Итак, когда доход в стандартной форме составляет $1,12\sigma$, то социальный статус оценивается как $1,00\sigma$, если доход равен – $1,56\sigma$, то социальный статус – $1,39\sigma$. Полная корреляция между двумя рядами данных имела бы место всякий раз, когда парными значениями были бы идентичные расстояния в долях σ от соответствующих средних. Например, если бы человек с оценкой по социологии, скажем, на $1,8\sigma$ выше среднего, располагал бы на

1,8 σ выше среднего значения для оценок по истории, то и все другие наблюдения в этих двух рядах были бы совершенно аналогичными. Поскольку r обратим, можно применить его к наблюдаемым значениям z_y и тем самым оценить соответствующую величину z_x .

В результате осуществленных вычислений, описанных выше, очевидность приобретает тот факт, что основная формула вычисления коэффициента корреляции Пирсона, хотя и вполне понятна, однако, мягко говоря, не очень удобна для вычисления «вручную», даже с использованием калькулятора. Поэтому существуют производные формулы – более громоздкие по виду, менее доступные осмыслению и не совсем понятные, однако существенно упрощающие расчеты. Кроме того, эти формулы, в свою очередь, дифференцируются на: «для сгруппированных данных» и «для несгруппированных данных». В данном издании мы не будем их приводить, так как «вручную» коэффициент корреляции вычисляется крайне редко (только в каких-то крайних случаях), а в основном для обработки реальных данных социологического исследования с целью определения корреляционных связей и зависимостей используются специальные компьютерные программы.

Основные свойства коэффициента корреляции Пирсона. Как мера тесноты и направления (линейной) связи между двумя признаками, коэффициент корреляции Пирсона r обладает следующими свойствами:

1) $r_{y/x} = r_{x/y}$ (поэтому, как правило, этот коэффициент обозначается просто одной буквой r);

2) чем ближе $r_{y/x}$ к +1 или -1, тем теснее связь. Чем ближе $r_{y/x}$ к 0, тем связь слабее. Если $r_{y/x} = 0$, то говорят, что связи нет. Для социальных процессов $r_{y/x}$ редко превышает |0,75|;

3) если $r_{y/x} > 0$, то связь прямая. Если $r_{y/x} < 0$, то связь обратная.

Сравнение r и r^2 . При поверхностном рассмотрении разница между коэффициентами r и r^2 кажется весьма тривиальной: налицо просто разные показатели степени, причем обе величины легко преобразуются друг в друга, как только одна из них вычислена.

Если рассмотреть такие преобразования на примере связи дохода респондентов и их оценки собственного социального статуса (см. *Табл. 4.16* и *4.17*) при коэффициенте корреляции Пирсона $r = 0,89$ получим коэффициент детерминации $r^2 = 0,79 \approx 0,8$ (что, кстати, вполне согласуется с результатами вычислений по нахождению r^2 , осуществленных с применением специальной формулы).

Тем не менее этой кажущейся незначительной разницей нельзя с легкостью пренебрегать, потому что оба коэффициента используются в двух различных, хотя и взаимосвязанных, аспектах ковариации.

Каждая из двух рассматриваемых величин механически выводится из другой посредством одной из двух процедур. Однако критерий r^2 , полученный из r , не дает полного представления о назначении и смысле объясняемой вариации, которая измеряется величиной r^2 . Понятие «квадрат наклона» ничего не говорит студенту. В такой формулировке, как «корень квадратный из объясняемой вариации», нельзя было бы усмотреть наклон. Хотя r и r^2 взаимозависимы, критерий r^2 измеряет всю ту долю полной вариации одной переменной, которая связана с другой или объясняется ею.

С другой стороны, критерий r оценивает динамический аспект этого отношения, измеряя скорость изменения одной переменной относительно другой, как это было показано в предыдущих примерах. Исходя из этого концептуального различия, можно утверждать, что критерий r является, главным образом, средством предсказания, например, ожидаемого уровня изменения одной переменной при наблюдаемом изменении другой. Как таковой он пригодился бы педагогам, психологам и другим исследователям, которых интересует единичное предсказание, а вот социологам – в значительно меньшей степени. С другой стороны, r^2 есть суммарная мера, взвешивающая влияние (или воздействие) одной переменной на другую.

Поскольку величина r представляет собой наклон, то она, очевидно, должна задавать направление преимущественно вверх или вниз в соответствии с тем, положительно или отрицательно связаны переменные. Это означает, что направление наклона отражает тип связи, который обозначается знаком «плюс» или «минус». Синоптическая мера r^2 , однако, не несет знака, ибо она выражает долю общей вариации.

Теперь ясно, что r и r^2 не являются взаимозаменяемыми; их также не следует непосредственно выводить друг из друга до тех пор, пока их структурный смысл не будет понят. Так как значение r всегда больше, чем значение r^2 , то случайное или преднамеренное использование r для выражения силы связи вводило бы читателя в заблуждение: например, могло бы показаться, что $r = 0,5$ означает довольно сильную связь, но $r^2 = 0,25$ показывает, что лишь 25% вариации каждой переменной связано с другой. Когда требуется подчеркнуть силу полной связи между двумя переменными, что часто бывает в социологических исследованиях, более подходящей статистикой является r^2 . Например, если выяснено, что 50% вариации в области преступности может быть связано с экономическим фактором, хотя кое-что остается еще необъясненным (какими факторами вызваны еще 50% вариации в области преступности), все же это является реальным продвижением на пути понимания явления, которое всего сто лет назад объясняли нечистой силой, наследственностью или свободной волей.

В итоге каждый критерий имеет свое собственное обозначение и соответствующее употребление. Тем не менее, критерий r пользуется большей популярностью по сравнению с r^2 . Возможно, отчасти это объясняется силой привычки, которая, вероятно, исчезнет, когда исследователи станут более чувствительными к нюансам количественного описания.

Делая общий вывод, подчеркнем, что используя коэффициент корреляции Пирсона, следует учитывать, что лучше всего он подходит для оценки взаимосвязи между двумя переменными, значения которых распределены нормально (т. е. согласно закону нормального распределения). Если распределение переменных

отличается от нормального, то этот коэффициент по-прежнему продолжает характеризовать степень взаимосвязи между признаками, однако в данном случае к нему уже нельзя применять методы проверки на значимость. Также коэффициент корреляции Пирсона не очень устойчив к выбросам (резко выдающимся значениям). То есть в тех случаях, когда есть резко выделяющиеся значения, можно ошибочно сделать вывод о наличии корреляции между переменными. Поэтому если распределение исследуемых переменных отличается от нормального или возможны выбросы, то лучше воспользоваться либо «непараметрическим аналогом», то есть коэффициентом ранговой корреляции Спирмена (о котором подробно говорилось в предыдущем подразделе), либо обратиться к другим методам определения корреляции, один из которых будет подробно описан далее.

4.9. Нелинейная регрессия.

Множественная и частная корреляция

Нелинейная регрессия. Всякий коэффициент корреляции показывает нам, в какой степени одну переменную можно предсказать или объяснить через другую. Из многих употребляемых в статистике показателей коэффициент r Пирсона является одной из наиболее распространенных, почти банальных мер. Однако применимость r основывается в основном на следующих двух условиях: 1) отношение между переменными линейно; 2) двумерное распределение гомоскедастично, то есть имеет постоянную условную дисперсию. Вообще говоря, оба эти условия выполняются, если маргинальные распределения нормальны. Поэтому одним из требований, предъявляемых к критерию r , часто является нормальность маргинальных распределений. Однако очень часто исходные данные приводят к нелинейным моделям. Это особенно справедливо для социологических наук, где многие распределения, такие как доход или размер семьи, сильно скошены, и линии регрессии, поэтому, в значительной мере нелинейны.

Так как критерий r определяет прямолинейную модель, то становится неподходящим в той мере, в какой связь между

переменными отклоняется от линейной. То же самое получается, если диаграмма рассеяния гетероскедастична, даже если при этом она остается линейной. Конечно, поскольку эти идеальные условия никогда не могут быть выполнены, исследователи вынуждены прибегать к аппроксимации. Но существует разумный предел, за который не следовало бы распространять аппроксимацию, если доступны более подходящие способы измерения.

Существуют различные методы, с помощью которых можно измерить криволинейные зависимости, но здесь будем рассматривать только один из них. Излагаемый метод есть метод корреляционного отношения, который часто обозначается через η (греческая строчная буква «эта»), но лучше выражается посредством η^2 . Хотя обычно «корреляционное отношение» отождествляют просто с η , во время обсуждения соответствующей темы, однако, станет ясно, что критерий η мало пригоден для практического применения; именно поэтому он является незначимым критерием корреляции. Поэтому термин «корреляционное отношение» будет использоваться как равнозначный с величиной η^2 .

Данный метод представляет собой относительно простую процедуру и связан с тем же принципом, на котором основано определение r^2 , а именно: на соотношении объясняемой и общей вариации. Вычисления в этом случае отличаются от определения r^2 в том отношении, что объясняемая вариация выводится из среднего значения по столбцам (строчкам) корреляционной таблицы, а не из гипотетической линии регрессии. Тем самым корреляционное отношение дополняет «пирсоновский» критерий r^2 , пользующийся популярностью и репутацией, вероятно, превышающими его практическую значимость в области социологических наук.

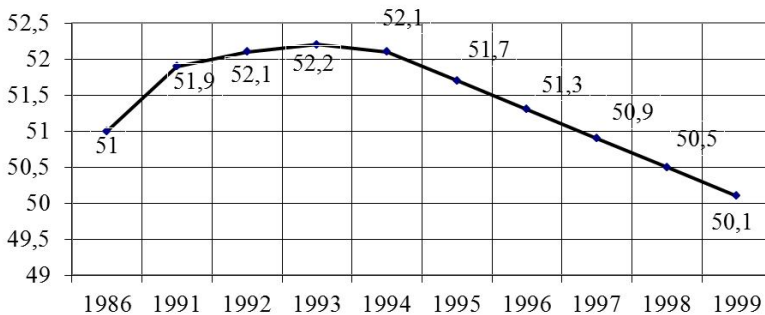
Вычисление корреляционного отношения. В целях простой иллюстрации вычислим корреляционное отношение между численностью населения Украины и годами с 1986 по 1999 по официальным данным статистики (см. *Табл. 4.18*)

Чтобы установить, соответствуют ли рассматриваемые данные линейной или нелинейной модели, надо, прежде всего, построить обычное корреляционное поле (см. *Рис. 4.7*), которое позволит определить это визуально.

Таблица 4.18

Численность населения Украины в 1986–1999 годах

Год	Численность (млн человек)
1986	51
1991	51,9
1992	52,1
1993	52,2
1994	52,1
1995	51,7
1996	51,3
1997	50,9
1998	50,5
1999	50,1

Рис. 4.7. Численность населения Украины в 1986–1999 годах:
корреляционное поле

Корреляционное поле, на которое наложена проведенная от руки линия тенденции, позволяет оценить силу и тип связи между двумя переменными. Можно было бы представить разброс и прямой линией, но кривая следует контуру данных более точно. Поэтому заключаем, что η^2 дает лучшее приближение, чем r^2 , и приступаем к вычислению η_{xy} .

Процедура вычисления значения η^2 совершенно аналогична той, что имела место при вычислении r^2 Пирсона. Необходимо вычислить полную вариацию и объясняемую вариацию, а затем найти отношение между ними. Соответствующая формула (для негруппированных данных) выглядит следующим образом:

$$\eta_{xy}^2 = \frac{\text{объясняемая_вариация}}{\text{полная_вариация}} = \frac{\sum (\hat{y}_i - \bar{y})^2}{\sum (y_i - \bar{y})^2},$$

где • k – число групп по факторному признаку;

- N – число единиц совокупности;
- $\hat{y}_i$ – внутригрупповое среднее значение признака Y ;
- $\bar{y}$ – общее среднее значение признака Y ;
- y_i – индивидуальные значения признака Y .

Данная формула представляет собой вычисление η^2 .

В случае, если данные сгруппированы, эта формула преобразовывается и имеет такой вид:

$$\eta_{yx}^2 = \frac{\text{объясняемая_вариация}}{\text{полная_вариация}} = \frac{\sum (\hat{y}_i - \bar{y})^2 \cdot n_j}{\sum (y_i - \bar{y})^2},$$

где • k – число групп по факторному признаку X ;

- N – число единиц совокупности;
- $\hat{y}_i$ – внутригрупповое среднее значение признака Y ;
- $\bar{y}$ – общее среднее значение признака Y ;
- y_i – индивидуальные значения признака Y ;
- n_j – частота значения признака в j -ой группе.

Процедуру в целом лучше всего описать по этапам, что мы и сделаем, начиная с таблицы 4.19.

1. Применяем обычный способ двумерной группировки данных, в нашем случае – это данные распределения численности населения Украины (признак Y) в 1986–1999 годах (признак X). Результаты группировки, приведены в таблице 4.19.

Таблица 4.19

Численность населения Украины в 1986–1999 годах

Год (x_i)		Численность (млн человек) (y_i)	$(y_i - \bar{y})^2$	$\bar{y}$ – среднее арифметическое – групповое	$(\hat{y}_i - \bar{y})^2$
1	1986	51	$(51 - 51,38)^2 = 0,1444$	$(51 + 51,9)/2 = 51,45$	$(51,45 - 51,38)^2 = 0,005$
2	1991	51,9	0,2704	51,45	0,005
3	1992	52,1	0,5184	52,133	0,567
4	1993	52,2	0,6724	52,133	0,567
5	1994	52,1	0,5184	52,133	0,567
6	1995	51,7	0,1024	51,5	0,014
7	1996	51,3	0,0064	51,5	0,014
8	1997	50,9	0,2304	50,5	0,77
9	1998	50,5	0,7744	50,5	0,77
10	1999	50,1	1,6384	50,5	0,77
N=10 Σ		513,8	4,876		4,049

2. Вычисляем общее среднее для значений признака Y по формуле:

$$\bar{y} = \frac{\sum y_i}{N} = \frac{513,8}{10} = 51,38.$$

3. Подсчитываем сумму квадратов отклонений по y , или общую вариацию по формуле:

$$\sigma^2 = \sum (y_i - \bar{y})^2 = 4,876.$$

4. Вычисляем среднее $\hat{y}$ для значений y в каждой группе, причем разделение на группы определяется исследователем с точки зрения здравого смысла. В нашем случае деление на группы определено – подчеркиванием.

5. Находим сумму квадратичных отклонений. Это и есть объясняемая вариация.

6. Делим объясняемую вариацию на полную вариацию и получаем, согласно определению, корреляционное отношение:

$$\eta_{yx}^2 = \frac{\text{объясняемая_вариация}}{\text{полная_вариация}} = \frac{4,049}{4,876} = 0,83.$$

Итак, было найдено, что 0,83% вариации численности населения Украины объясняется изменением года. Действуя точно в той же последовательности, но уже со средними для строк, а не для столбцов, можно определить η_{xy}^2 – часть вариации x , объясняемую через y . Тот факт, что эти два значения η^2 различны, указывает, что, в противоположность r^2 , критерий η^2 не является обратимым. В общем случае относительная вариация в строках и столбцах не совпадает в точности, поэтому значения η_{xy}^2 чаще всего не равно значению η_{yx}^2 . В крайнем, весьма маловероятном случае, вариация по столбцам может отсутствовать и вовсе, оставаясь существенной по строкам, так что значение η_{yx}^2 было бы равно единице, а значение η_{xy}^2 обнаруживало бы лишь умеренную степень связи.

Сравнение статистических показателей r^2 и η^2 :

$r^2 = 0$, если x и y независимы.

$r^2 = \eta_{x/y}^2 = 1$, тогда и только тогда, когда имеется строгая линейная связь между x и y .

$r^2 \leq \eta_{x/y}^2$, тогда и только тогда, когда имеется строгая нелинейная функциональная зависимость x и y .

$r^2 = \eta_{x/y}^2 < 1$, тогда и только тогда, когда регрессия x и y строго линейная, но нет функциональной зависимости.

$r^2 < \eta_{x/y}^2 < 1$ означает, что нет функциональной зависимости и существует нелинейная кривая регрессии.

Условия применимости критерия η^2 . Границы применения критерия η^2 более широки, нежели для критерия r^2 Пирсона. Так, не накладывается никаких ограничений на форму маргинальных распределений, и они могут быть нормальными, скошенными или даже бимодальными; линия регрессии может иметь любую криволинейную форму. Как переменную по одной оси можно использовать даже качественную переменную (по двум осям –

нельзя, ибо тогда невозможно было бы найти среднее арифметическое и пришлось бы вычислять коэффициент взаимной сопряженности C). Можно предложить только один ограничивающий фактор: совокупность нанесенных на диаграмму точек должна в общем случае быть гомоскедастична, по крайней мере, в одном направлении; это необходимо по той же причине, что и для r Пирсона, а именно потому, что синоптическая мера обычно более значима, если ее получают из однородных данных. Но это – скорее логическое, а не математическое ограничение. По этой же причине иногда нецелесообразно применять среднее для данных с большой дисперсией.

Принципы интерпретации. Так как гипотетическая линия регрессии меняет направление из-за своей кривизны и изгибов, то весьма трудно интерпретировать величину η^2 тем же способом, как и r^2 : «чем выше значение x , тем выше значение y » (методически: «среднее изменение y , приходящееся на единицу изменения x »). Такая интерпретация была бы лишена смысла, поскольку линия регрессии может менять направление, следовательно, величина не эквивалентна постоянному наклону r . Действительно, η не выполняет никакой определенной функции в измерении корреляции; она лишь является корнем квадратным из величины η^2 , которая больше подходит для нелинейной корреляции.

Поскольку всякая совокупность нанесенных на диаграмму точек всегда будет содержать в себе определенную нелинейность, то значение η^2 , отражающее криволинейность и линейность одинаково хорошо, всегда будет более точной оценкой для объясняемой вариации и, следовательно, будет всегда выше, чем значение r^2 . Критерий r^2 всегда предельно расширяет возможности статистического объяснения, ибо он автоматически приспособливается к конфигурации данных. С другой стороны, критерий r^2 является более жестким, что ограничивает его возможности отражением лишь линейной модели. В случае полной линейной зависимости $r^2 = \eta^2$.

Причина, по которой критерий r^2 может более полно отражать объясняемую вариацию, состоит просто-напросто в том, что он

делит искривленную линию регрессии на сегменты, которые всегда находятся в большей близости к точкам совокупности, чем одна прямая линия. Поэтому остатки будут меньше. В строго линейной совокупности с одним только направлением такая гибкость, разумеется, не нужна.

На рисунке 4.7 видно, что для данной совокупности подходит не прямая, а именно кривая линия регрессии. Однако в математических методах нет ничего, что исключало бы применение линейной формулы к нелинейным данным. При этом можно было бы уловить в какой-то мере линейность, которая в действительности может вообще отсутствовать. Кривая линия просто лучше приближает данные и поэтому больше подходит для целей исследования.

Большая точность криволинейного приближения становится все очевиднее, когда используется линия регрессии по ее назначению, а именно для предсказания одной переменной по другой. Очевидно, значения y , определяемые по прямой линии, приведут к значительно большим ошибкам, чем в случае использования кривой, ибо прямая линия становится все менее и менее представительной по мере того, как данные становятся все более и более нелинейными. В самом деле, на краях распределения данные далеко отклоняются от прямой линии, так что предсказуемые значения сильно отклоняются от наблюдаемых.

Более того, если к нелинейной зависимости применяется линейная модель, нормирующая способность величины r^2 уменьшается, потому что индекс не может меняться от нуля до единицы. Значения r^2 , если его ошибочно применять к нелинейным данным, не достигало бы единицы. Но критерий η^2 , как это уже было показано, не имеет такого ограничения и теоретически может достигать единицы. В этом случае предполагается, что при определенных условиях, когда r^2 не подходит, то обращаются к η^2 .

Предосторожности в применении критерия η^2 . Чтобы использовать формулу для η^2 , необходимо сгруппировать данные по сегментам, а такая группировка уже влечет за собой в определенной степени произвол. Доводя этот процесс до абсурда,

можно выделить столько групп классификации, или интервалов группировки, сколько имеется точек, располагая каждую наблюдаемую точку в отдельном столбце. Кривая регрессии проходила бы тогда через другую точку. При такой группировке значения η^2 было бы в точности равно единице, ибо всякое отклонение внутри строк исключалось бы. С другой стороны, интервалы группировки можно было бы сделать абсурдно большими, и тогда не удалось бы выявить форму распределения. В таком случае остатки были бы слишком велики для того, чтобы судить о кривизне распределения, и мы вновь столкнулись бы с известным явлением неправильной группировки.

Виды нелинейной формы связи. В случае, когда $r^2 \leq \eta_{x/y}^2$, между признаками существует строгая нелинейная связь. Задача социолога отыскать уравнение функциональной зависимости, определяющее эту связь, то есть уравнение регрессии. В данном издании представим без особых объяснений только два уравнения криволинейной зависимости: параболическое и гиперболическое.

$\bar{y}_x = a + b \cdot x + c \cdot x^2$ – уравнение, определяющее параболу.

Для нахождения параметров a , b , c необходимо решить следующую систему уравнений:

$$\begin{cases} \sum y = na + b \sum x + c \sum x^2 \\ \sum xy = a \sum x + b \sum x^2 + c \sum x^3 \\ \sum yx^2 = a \sum x^2 + b \sum x^3 + c \sum x^4. \end{cases}$$

Для сгруппированных данных формулы поиска параметров a , b , c выглядят так:

$$\begin{cases} \sum y n_y = na + c \sum (x - \bar{x})^2 \cdot n_x \\ \sum y(x - \bar{x}) \cdot n_y \cdot n_x = b \sum (x - \bar{x})^2 n_x \\ \sum y(x - \bar{x})^2 \cdot n_y \cdot n_x = a \sum (x - \bar{x})^2 \cdot n_x + c \sum (x - \bar{x})^4 \cdot n_x. \end{cases}$$

В случае если корреляционное поле представляет собой кривую, которая может быть описана гиперболой, применяется

формула: $y_x = a + \frac{b}{x}$.

Для нахождения параметров a , b необходимо решить систему уравнений:

$$\begin{cases} \sum y/x = a \sum 1/x + b \sum 1/x^2 \\ \sum y = an + b \sum 1/x \end{cases} \quad \text{– для несгруппированных данных;}$$

$$\begin{cases} \sum y/x \cdot n_n \cdot n_y = a \sum n_x/x + b \sum n_x/x^2 \\ \sum y \cdot n_y = an + b \sum n_x/x \end{cases} \quad \begin{array}{l} \text{– для сгруппированных} \\ \text{данных.} \end{array}$$

Частная и множественная регрессия и корреляция. Вывод о связи может быть сделан только на основании анализа всей совокупности связей в системе «изучаемый процесс – факторные признаки», поэтому рассчитывают не отдельные коэффициенты, а таблицу коэффициентов.

Так, если S_0 – показатель, отражающий уровень интереса студентов к социологии, а F_1, F_2, F_3 – факторы учебного процесса, отражающие содержание программы, уровень квалификации преподавателей и объем программы в часах, матрица коэффициентов связи может выглядеть так:

Таблица 4.20

Матрица коэффициентов связи (эмпирический пример)

	S_0	F_1	F_2	F_3
S_0	1	0,89	0,68	0,75
F_1	0,89	1	0,56	0,95
F_2	0,68	0,56	1	0,11
F_3	0,75	0,95	0,11	1

Чисто внешне значимы все три фактора, однако при анализе внутренних связей можно заметить, что оценка содержания программы оказалась зависимой от ее объема в часах, поэтому, несмотря на то что уровень связи изучаемого процесса с показателем «содержание программы» выше, детерминирующим следует считать фактор «объем программы в часах», так как именно он определяет меру содержательности программы. Кроме того, на содержание программы оказывает значимое влияние и фактор квалификации преподавателей.

Построение математической факторной модели включает оценку количественного влияния факторов на изучаемый процесс. Модель разрабатывается после качественного анализа влияния факторов и требует включения только тех факторов, влияние которых на изучаемый процесс доказано на предыдущем этапе. Модель, как правило, представляется в виде регрессионной функции вида: $y = f(x_1, x_2, \dots, x_k)$. Вид функции выбирается исходя из качественного анализа процесса или подбирается путем перебора. Для моделирования вида функции часто используют стандартные пакеты типа Statgraf или Statistica для Windows.

Ранее было показано, как можно по опытным данным найти зависимость одной переменной от другой, а именно, как построить уравнение регрессии вида $y = F(x_1)$. Если исследователь изучает влияние нескольких переменных $x_1, x_2, \dots, x_k$, на результирующий признак y , то возникает необходимость в умении строить регрессионное уравнение более общего вида, то есть $y = a + b_1x_1 + b_2x_2 + \dots + b_kx_k$ – уравнение множественной регрессии, где $a, b_1, b_2, \dots, b_k$ – постоянные коэффициенты, называемые частными коэффициентами регрессии.

В связи с последним уравнением необходимо рассмотреть следующие вопросы: а) как по эмпирическим данным вычислить коэффициенты регрессии $a, b_1, b_2, \dots, b_k$; б) какую интерпретацию можно приписать этим коэффициентам; в) как оценить тесноту связи между y и каждым из x_i в отдельности (при элиминировании действия остальных); г) оценить тесноту связи между y и всеми переменными $x_1, x_2, \dots, x_k$ в совокупности.

Рассмотрим этот вопрос на примере построения двухфакторного регрессионного уравнения. Предположим, что изучается зависимость недельного бюджета свободного времени (y) от уровня образования (x_1) и возраста (x_2) определенной группы трудящихся по данным выборочного исследования. Будем искать эту зависимость, обратившись к линейному уравнению следующего вида: $y = a + b_1x_1 + b_2x_2$.

При расчете коэффициентов уравнения множественной регрессии полезно преобразовать исходные эмпирические дан-

ные следующим образом: $z_{ij} = \frac{x_{ij} - \bar{x}_i}{\sigma_i}$, при этом уравнение

множественной регрессии примет вид $y = c_1z_1 + c_2z_2$, где c_1 и c_2 находятся из следующей системы уравнений:

$$\begin{cases} c_1 + r_{12}c_2 = r_{1y} \\ c_1r_{12} + c_2 = r_{2y} \end{cases}$$

решая которую, получаем, что $c_1 = \frac{r_{1y} - r_{12}r_{2y}}{1 - r_{12}^2}$, $c_2 = \frac{r_{2y} - r_{12}r_{1y}}{1 - r_{12}^2}$,

где r – коэффициент парной корреляции между признаками.

c_1 и c_2 называются стандартизированными коэффициентами регрессии. Следовательно, зная коэффициенты корреляции между изучаемыми признаками, можно подсчитать коэффициенты регрессии. Подставим конкретные значения всех возможных r из таблицы 4.21. Пусть в результате исследования N человек получены эмпирические значения, сведенные в следующую таблицу (в каждом столбце представлены несгруппированные данные).

$$\text{Тогда } c_1 = \frac{0.556 - (-0.131)(-0.027)}{1 - (-0.027)^2} = 0.55.$$

Аналогично $c_2 = -0,12$, и уравнение регрессии запишется в виде $y = 0,55z_1 - 0,12z_2$.

Таблица 4.21

Матрица коэффициентов связи (теоретический пример)

	Y	X ₁	X ₂	Z ₁	Z ₂
1	y ₁	x ₁ ¹	x ₁ ²		
2	y ₂	x ₂ ¹	x ₂ ²		
3	y ₃	x ₃ ¹	x ₃ ²		
...	...	...	...		
n	y _n	x _n ¹	x _n ²		
	$\bar{y}_i$	$\bar{x}_1$	$\bar{x}_2$		
	σ_y	σ_{x1}	σ_{x2}		

r_{12}	r_{1y}	r_{2y}
-0.027	0.556	-0.131

Коэффициенты $a, b_1, b_2, \dots, b_k$ исходного регрессионного уравнения находятся по формулам:

$$b_1 = c_1 \frac{\sigma_y}{\sigma_1}; \quad b_2 = c_2 \frac{\sigma_y}{\sigma_2};$$

$$a = \bar{y} - b_1 x_1 - b_2 x_2.$$

Подставляя в эти формулы полученные данные, будем иметь:

$$b_1 = c_1 \frac{\sigma_y}{\sigma_1} = 3,13; \quad b_2 = c_2 \frac{\sigma_y}{\sigma_2} = 0,17;$$

$$a = \bar{y} - b_1 x_1 - b_2 x_2 = 8,56.$$

Как же следует интерпретировать это уравнение? Например, значение b_2 показывает, что в среднем недельный бюджет свободного времени при увеличении возраста на один год и при фиксированном значении x_1 уменьшается на 0,17 часа. Аналогично интерпретируется b_1 .

Коэффициенты b_1, b_2 можно в то же время рассматривать и как показатели тесноты связи между переменными y и, например x_1 , при постоянстве x_2 .

Аналогичную интерпретацию можно применять и к стандартизированным коэффициентам регрессии c_1 и c_2 . Однако поскольку

ку c_1 и c_2 вычисляются исходя из нормированных переменных, они являются безразмерными и позволяют сравнивать тесноту связи между переменными, измеряемыми в различных единицах. Например, в вышеприведенном примере x_1 измеряется в классах, а x_2 – в годах. c_1 и c_2 позволяют сравнить, насколько x_1 теснее связан с y , чем x_2 .

Частный коэффициент корреляции – показатель, который характеризует тесноту и направление связи между результирующим признаком (y) и факторным признаком (x_i) при элиминировании остальных признаков.

Частный коэффициент корреляции записывается $r_{y1.2}$ и вычисляется по следующей формуле:

$$r_{y1.2} = \frac{r_{y1} - r_{y2}r_{12}}{\sqrt{(1 - r_{y2}^2)(1 - r_{12}^2)}}.$$

Для характеристики степени связи результирующего признака y с совокупностью независимых переменных служит множественный коэффициент корреляции $R_{y(1...k)}^2$, который вычисляется по формуле (иногда он выражается в процентах):

$$1 - R_{y(1...k)}^2 = (1 - r_{y1}^2)(1 - r_{y2.1}^2) \dots (1 - r_{yn.23...(k-1)}^2).$$

Так, для вышеприведенного примера множественный коэффициент корреляции равен:

$$R_{y(12)}^2 = 1 - (1 - r_{y1}^2)(1 - r_{y2.1}^2) = 1 - (1 - 0,56^2)(1 - 0,140^2) = 0,323 = 32\%.$$

Множественный коэффициент корреляции показывает, что включение признаков x_1 и x_2 в уравнение $y = 8,35 + 3,14x_1 - 0,166x_2$ на 32% объясняет изменчивость результирующего фактора. Чем больше $R_{y(1...k)}^2$, тем полнее независимые переменные $x_1, x_2, \dots, x_k$ описывают признак y . Обычно $R_{y(1...k)}^2$ служит критерием включения или исключения новой переменной в регрессионное уравнение. Если $R_{y(1...k)}^2$ мало изменяется при включении новой переменной в уравнение, то такая переменная отбрасывается.

Вопросы и задания для самоконтроля

1. Дайте определения: взаимная сопряженность, совместное появление, ковариация.

2. Охарактеризуйте следующие связи как простые или сложные и дайте свое обоснование: продолжительность брака и размер семьи; доход и социальный статус; смертный приговор и уменьшение случаев убийства; психическое расстройство и самоубийство; диаметр и длина окружности.

3. Дайте определения: корреляционное поле, прямолинейность, криволинейность, корреляционная таблица, маргинальные распределения, вариабельность, двумерная совокупность данных.

4. Если маргинальные распределения нормальны по форме, будет ли рассеяние точек обязательно гомоскедастичным?

5. Дайте определения: таблица 2×1 , таблица 2×2 , статистическая связь, односторонняя и двусторонняя связь, ожидаемая и наблюдаемая частота, маргинал, распределение совместных частот.

6. При исследовании психически ненормальных людей 94% продемонстрировали повышенную конфликтность в поведении перед началом психического заболевания: А). Доказывает ли это наблюдение связь между состоянием конфликта и психическим заболеванием, и почему Вы думаете именно так? Б). Какой бы Вы сделали вывод, если бы 94% из группы «нормальных» людей так же переживали состояние конфликта? В). Если бы 50% из группы «нормальных» людей переживали состояние конфликта, какой бы вывод вы сделали? Г). Достаточно ли иметь таблицу 2×1 для того, чтобы доказать наличие связи?

7. Дайте определения: ранг, ранговый порядок, равные интервалы, порядковые числа, корреляция ранговых последовательностей.

8. При каких обстоятельствах ранжируются качественные данные?

9. Приведите пример уменьшения корреляции в случае объединения рангов.

10. Опишите алгоритм вычисления коэффициента Спирмена.
11. Каков принцип вычисления множественного коэффициента корреляции для ранговых рядов?
12. Дайте определения: коэффициент детерминации, объясняемая вариация, коэффициент недетерминированности, необъясняемая вариация, линия регрессии, линия наименьших квадратов, коэффициент корреляции моментов произведений, угол наклона, общее отклонение, необъясняемое отклонение.
13. Как следует понимать необъясняемые остатки: как результат случая либо как результат воздействия определяющих факторов?
14. Предположим, что $r=0,3$ для связи между школьными оценками и часами подготовки к занятиям. Проанализируйте эту «низкую корреляцию». Действительно ли подготовка не влияет на оценки?
15. Для использования требуется, чтобы рассеяние наблюдений относительно линии регрессии было гомоскедастичным. Объясните, почему?
16. Допустимо ли вычисление η^2 , если линия регрессии плавная?
17. Можно ли вычислить η^2 для двух качественных переменных? Почему?
18. Как порядок столбцов (строк) влияет на числовое значение η^2 ?
19. При каких табличных условиях числовые значения r^2 и η^2 совпадают?
20. Покажите графически условия, при которых значение r^2 приблизительно равнялось бы нулю, а η^2 – единице.
21. Проверьте графически, что значение η^2 можно сделать близким к единице, используя столько же столбцов, сколько имеется величин.
22. Подсчитайте η_{yx}^2 и η_{xy}^2 для динамики сети высших учебных заведений Украины III–IV уровня аккредитации и численности студентов в них (см. Таблица 4.13, с. 170).
23. Что фиксирует частный коэффициент корреляции?

Раздел V. ПРИМЕНЕНИЕ МАТЕМАТИЧЕСКИХ МЕТОДОВ ПРИ ИССЛЕДОВАНИИ МАЛЫХ СОЦИАЛЬНЫХ ГРУПП

5.1. Социометрический опрос как метод социологического исследования малых социальных групп: специфика и общая схема действий

Математика – фундаментальная наука, предоставляющая общие языковые средства другим наукам. С помощью этих средств становится возможным определение некоторых исчисляемых свойств объекта, находящегося в фокусе исследования и поддающегося измерению, уточнение его структуры, порядка составных частей, элементов и специфики отношений между этими элементами. Другими словами, с помощью математики, тех или иных математических методов, процедур становится возможным выявление структурных взаимосвязей и выведения общих законов, обуславливающих эти взаимосвязи.

Все математические процедуры, детально рассмотренные нами в предыдущих разделах данного издания, связанные с агрегированием данных, (определением мер средней тенденции и вариации) и тем более с определением взаимосвязи между признаками (вычислением коэффициентов корреляции и т. п.), в социологии применяются преимущественно в случае исследования больших социальных групп.

Если среднее арифметическое, моду, медиану, дисперсию и некоторые другие меры усреднения можно вычислить и для относительно небольшого числа единиц исследуемой совокупности, то коэффициенты корреляции и вовсе значимы только при условии большого объема этой совокупности. Однако, исследования малых социальных групп, то есть те, которые касаются микроуровня социального взаимодействия, могут принести не менее важные, с теоретической и практической точек зрения, результаты.

Исследование малых групп имеет в качестве своей предпосылки характеристику некоторой «статики» группы: определение ее границ, состава, композиции. Но, естественно, что главной задачей социально-психологического анализа является изучение процессов, которые происходят в жизни группы. Рассмотрение их важно с двух точек зрения: во-первых, необходимо выяснить, как общие закономерности общения и взаимодействия реализуются именно в малой группе, потому что здесь создается конкретная ткань коммуникативных, интерактивных и перцептивных процессов; во-вторых, нужно показать, каков механизм, посредством которого малая группа «доводит» до личности всю систему общественных влияний, в частности, содержание тех ценностей, норм, установок, которое доминируют и просто имеют место в больших группах. Вместе с тем важно выявить и обратное движение, а именно: каким образом активность личности в группе реализует усвоенные влияния и осуществляет определенную отдачу. Значит, важно дать срез того, что происходит в малых группах. Но это только один аспект проблемы. Другая не менее важная задача состоит в том, чтобы показать, как развивается группа, какие этапы развития она проходит, как модифицируются на каждом из этапов различные групповые процессы [4; 8; 18].

Общеизвестно, что одним из эффективных способов изучения отношений в малой группе является метод социометрии. Вопрос изучения малых групп при помощи этого метода является актуальным для современной социологии и социальной психологии.

В 30-е годы нашего столетия Я. (Дж.) Л. Морено предложил термин «социометрия», а также разработал особую социо-психологическую теорию, согласно которой изменение психологических отношений в малой группе является главным условием изменений во всей социальной системе [7; 16].

Термин «социометрия» возник в конце XIX века, в связи с описанием возможных способов измерения социального влияния одних социальных групп на другие. Я. Морено, как основатель социометрии дает довольно убедительное ее теоретическое и идеологическое обоснование как метода познания и измерения

социальных явлений. Во-первых, социометрия есть общая теория социальных групп; во-вторых, социометрия означает всякое измерение всех социальных отношений; в-третьих, социометрией называется математическое изучение психологических свойств населения, экспериментальная техника и результаты, полученные при применении количественного и качественного методов. Объектом социометрической теории являются реально существующие малые социальные группы, имеющие достаточный опыт совместной групповой жизни. Предметная область социометрии – эмоциональные отношения людей в группах (симпатии, неприязнь, безразличие). Созданная на основе взглядов К. Маркса, О. Конта и З. Фрейда социометрия противопоставляется как бихевиоризму, наблюдающему лишь внешне поведение людей, так и фрейдизму с его акцентом на внутренних, глубинных процессах человеческого поведения. По Я. Морено, эмоциональные отношения людей в группах представляют атомистическую структуру общества, которая недоступна простому наблюдению и может быть вскрыта только с помощью, так называемой, социальной микроскопии. «Микросоциология, – писал Морено, – фактически возникла с появлением моей теории «социальной микроскопии». В соединении с социометрическими приемами она положила начало теоретическим и практическим основам микросоциологии...». Изучение первичных «атомистических структур» человеческих отношений рассматривалось Морено как предварительная и необходимая основная работа для большинства макросоциологических исследований. Одно из центральных понятий этой теории – «теле» – термин, означающий простейшую единицу чувства, передаваемую от одного индивида к другому, детерминирующую количество и успешность межличностных отношений, в которые они вступают [16].

Суть «общей теории социометрии» состоит в утверждении того, что социальные системы являются притягательно/отталкивающе/нейтральными системами, включающими в себя не только объективные, внешне проявляемые отношения (макρο-структура), но и субъективные, эмоциональные отношения, часто невидимые внешне (микроструктура). Цель социометрической

теории – сформулировать законы эмоциональных отношений в группах.

Основные положения теории Я. Морено:

1) социальный атом общества – это не отдельный индивид, а их взаимодействие;

2) закон социальной гравитации: сплоченность группы прямо пропорциональна влечению участников друг к другу;

3) социологический закон: высшие формы коллективной организации развиваются из простых форм;

4) социодинамический закон: внутри некоторых групп человеческие привязанности распространяются неравномерно.

В целом теория Я. Морено довольно активно критиковалась, особенно со стороны психологов и социологов, например, Л. Десева, В. Ядова и др. И по большей части эта критика была вполне правомерной, так как многие центральные идеи Я. Л. Морено представляются довольно утопичными. В то же время он изобрел метод, который оказался в социологии и социальной психологии чрезвычайно эффективным. На сегодняшний день социометрия в качестве системы прикладных методов для изучения отношений в малых группах нашла широкое применение в исследовательской практике зарубежных и отечественных ученых. Внедрение этого метода в современную науку связано с именами таких известных ученых, как Я. Л. Коломинский, Е. С. Кузьмин, В. И. Паниотто, В. А. Ядов и др.

Социометрическая техника, разработанная Я. Морено, применяется для диагностики межличностных и межгрупповых отношений в целях их изменения, улучшения и совершенствования. С помощью социометрии можно изучать типологию социального поведения людей в условиях групповой деятельности, судить о социально-психологической совместимости членов конкретных групп.

Общая схема действий при социометрическом исследовании заключается в следующем. После постановки задач исследования и выбора объекта измерения формулируются основные гипотезы и положения, касающиеся возможных критериев опроса членов групп. Здесь не может быть полной анонимности, иначе

социометрия окажется малоэффективной. Требование экспериментатора раскрыть свои симпатии и антипатии нередко вызывает внутренние затруднения у опрашиваемых и проявляется у некоторых людей в нежелании участвовать в опросе. Когда вопросы или критерии социометрии выбраны, они заносятся на специальную карточку или предлагаются в устном виде по типу интервью. Каждый член группы обязан отвечать на них, выбирая тех или иных членов группы в зависимости от большей или меньшей склонности, предпочтительности их по сравнению с другими, симпатий или, наоборот, антипатий, доверия или недоверия и т. д. Членам группы предлагается ответить на вопросы, которые дают возможность обнаружить их симпатии и антипатии друг к другу, выявить лидеров и тех членов группы, которых группа не принимает. Исследователь зачитывает вопросы и дает респондентам следующую инструкцию: «Напишите на бумажках под цифрой «1» фамилию члена группы, которого Вы выбрали бы в первую очередь, под цифрой «2» – кого бы Вы выбрали, если бы не было первого, под цифрой 3 – кого бы Вы выбрали, если бы не было первого и второй». После того, как исследователь зачитает вопросы о личных отношениях, он проводит краткий инструктаж о правилах проведения исследования, которых необходимо придерживаться. С целью подтверждения достоверности ответов исследование может проводиться в группе несколько раз. Для повторного исследования берутся другие вопросы.

Примеры вопросов для изучения деловых отношений:

1. а) кого из членов «своей» группы Вы попросили бы, в случае необходимости, предоставить помощь в подготовке к занятиям (в первую, вторую, третью очередь)?

б) кого из группы Вы не хотели бы просить, в случае необходимости, предоставлять Вам помощь в подготовке к занятиям?

2. а) с кем бы Вы поехали в продолжительную служебную командировку?

б) кого из членов своей группы Вы не взяли бы в служебную командировку?

3. а) кто из членов группы лучше справится с обязанностями лидера (старосты и т.д.)?

б) кому из членов группы тяжело будет исполнять обязанности лидера?

Примеры вопросов для изучения личных отношений:

1. а) к кому в своей группе Вы обратились бы за советом в трудной жизненной ситуации?

б) с кем из группы Вам не хотелось бы ни о чем советоваться?

2. а) если бы все члены Вашей группы жили в общежитии, с кем из них Вам хотелось бы поселиться в одной комнате?

б) если бы всю Вашу группу переформировали, кого из ее членов Вы не хотели бы оставить в своей группе?

3. а) кого из группы Вы пригласили бы на День рождения?

б) кого из группы Вы не хотели бы видеть на своем Дне рождения?

При этом социометрическая процедура может проводиться в двух формах: непараметрической и параметрической. В первом случае респонденту предлагается ответить на вопросы социометрической карточки без ограничения числа выборов. Достоинством непараметрической процедуры является то, что она позволяет выявить так называемую эмоциональную экспансивность каждого члена группы, сделать срез многообразия межличностных связей в групповой структуре. Однако при увеличении размеров группы до 12–16 человек этих связей становится так много, что без применения специальной вычислительной техники проанализировать их становится весьма трудно. Другим недостатком непараметрической процедуры является большая вероятность получения случайного выбора (при параметрической процедуре это число ограничивается, в зависимости от общей численности группы и, соответственно, от вероятности каждого члена группы быть выбранным). Осознание этих недостатков привело исследователей к необходимости изменения процедуры опроса таким образом, чтобы снизить вероятность случайного выбора. Так родился второй вариант – параметрическая процедура социометрического опроса, с ограничением

числа выборов. При этом респондентам предлагается выбирать строго фиксированное число из всех членов группы. Например, в группе из 25 человек каждому предлагают выбрать лишь 4 или 5 человек. Недостатком этой процедуры является невозможность раскрыть многообразие взаимоотношений в группе. Представляется возможным выявление только наиболее субъективно значимых связей. Социометрическая структура группы в результате такого подхода будет отражать лишь наиболее типичные, «избранные» коммуникации, отсутствует возможность сделать выводы относительно эмоциональной экспансивности членов группы.

Социометрическая карточка, или социометрическая анкета, составляется на заключительном этапе разработки программы. В ней каждый член группы должен указать свое отношение к другим членам группы по ряду критериев (например, с точки зрения совместной работы, участия в решении деловой задачи, проведения досуга, в игре и т. д.) Критерии определяются в зависимости от программы данного исследования: изучаются ли отношения в производственной группе, группе досуга, во временной или стабильной группе [3; 11–13].

Использование социометрии позволяет проводить измерение авторитета формального и неформального лидеров для перегруппировки людей в командах так, чтобы снизить напряженность в коллективе, возникающую из-за взаимной неприязни некоторых членов группы. Социометрическая методика весьма полезна в прикладных исследованиях, особенно в работах по совершенствованию отношений в коллективе. Однако она не является радикальным способом разрешения внутригрупповых проблем, причины которых следует искать не в симпатиях и антипатиях членов группы, а в более глубоких источниках. Надежность процедуры зависит, прежде всего, от правильного отбора критериев социометрии, что диктуется программой исследования и предварительным знакомством со спецификой группы.

Резюмируя и обобщая сказанное, подчеркнем, что *социометрическая процедура может иметь целью:*

- а) изучение внутригрупповой структуры: выявление «социо-

метрических позиций», то есть относительного авторитета членов группы по признакам симпатии-антипатии, где на крайних полюсах оказываются «лидер» группы и «отвергнутый» (с помощью вычисления персональных социометрических индексов – «П.С.И.», в том числе, на основе социоматрицы);

б) изучение характера отношений в группе: степени сплоченности-разобщенности в группе (с помощью вычисления групповых социометрических индексов – «Г.С.И.»);

в) обнаружение внутригрупповых подсистем, сплоченных образований, во главе которых могут быть свои неформальные лидеры (на основе анализа и сопоставления «Г.С.И.» и «П.С.И.», в том числе, с помощью построения социограмм).

Более подробно о реализации каждой из представленных выше целей речь пойдет соответственно в подразделах 5.3, 5.4 и 5.5 данного издания.

5.2. Общая процедура проведения социометрического исследования

В одной из наиболее удачных, на наш взгляд, версий описания особенностей и специфики социометрии, представленной в одном из электронных научных источников под названием «Энциклопедия психодиагностики», сама процедура социометрического исследования подразделяется на шесть последовательных этапов [15]. Опишем более подробно эти этапы.

Этап I включает в себя **разработку программы** исследования: определение объекта, предмета, постановку целей и задач, формулировку гипотез.

Этап II связан с разработкой **социометрических критериев**. Социометрические критерии формируются в виде вопросов, ответы на которые служат основанием для установления отношений в группе. Типы критериев:

- формальные/неформальные (формальные критерии направлены на изучение взаимоотношений в ведущем виде деятельности группы; неформальные – на изучение неформальных отношений);

- двойные/одинарные (двойные – изучают отношения партнерства; одинарные – отношения лидерства и подчинения);
- прогностические – позволяют измерить адекватность отражения реальной картины взаимоотношений в сознании конкретного индивида;
- сильные значимые/слабые незначимые (сильные значимые критерии показывают наиболее глубокие, стабильные отношения; слабые незначимые – затрагивают только лишь поверхностные отношения).

Требования к критериям:

- а) ограниченное количество (3–4);
- б) логическая взаимосвязанность;
- в) критерии должны вызывать интерес у респондентов;
- г) смысл критериев должен быть (одинаково) понятен респондентам;
- д) критерии должны быть сформулированы максимально конкретно, лучше на примере реальной ситуации.

Этап III предполагает **выбор типа процедуры** социометрического исследования. Известны два основных типа: параметрическая процедура и непараметрическая процедура.

Параметрическая процедура была разработана для снижения вероятности случайного выбора. Ее суть состоит в ограничении количества выборов каждого их участников (обычно, для групп в 22–25 участников минимальная величина «социометрического ограничения» должна варьироваться в пределах 4–5 выборов). Величина ограничения (d) – называется социометрическим ограничением или лимитом выборов. С помощью введения этой величины можно стандартизировать внешние условия выборов в группах разной численности, то есть задать им определенные параметры. Величина d находится исходя из одинаковой для всех групп вероятности случайного выбора. Формулу определения такой вероятности предложили в свое время Я. Морено и Е. Дженнингс:

$$P(A) = \frac{d}{(N - 1)},$$

где $\bullet P(A)$ – вероятность случайного события социометрического выбора, которая обычно находится в пределах $[0,2-0,3]$;

- N – число членов группы.

Достоинства параметрической процедуры:

- повышает надежность полученных данных;
- облегчает статистическую обработку;
- позволяет стандартизировать условия выбора в группах различной численности;
- как следствие, появляется возможность сравнивать взаимоотношения в разных группах.

Недостатки параметрической процедуры:

- не раскрывает всего многообразия взаимоотношений в группе;
- социометрическая структура группы в результате такого подхода будет отражать лишь наиболее типичные, «избранные» коммуникации;
- введение «социометрического ограничения» не позволяет судить об эмоциональной экспансивности членов группы.

Непараметрическая процедура не предполагает каких-либо ограничений в количестве выборов, которые могут быть сделаны респондентами. Каждый респондент может предпочесть либо отвергнуть любое количество членов группы, кроме себя самого. Следовательно, по любому из предложенных критериев может быть выбрано $(N - 1)$ человек, что представляет собой основную константу социометрии (где N – количество человек в группе). Для непараметрической процедуры эта константа всегда одинакова как для человека, делающего выбор, так и для человека, выбор получающего.

Достоинства непараметрической процедуры:

- позволяет выявить эмоциональную экспансивность каждого члена группы;
- при использовании данного типа процедуры делается срез всего многообразия связей в групповой структуре.

Недостатки непараметрической процедуры:

- при ручной обработке результатов социометрии измерения возможны только в небольших группах (до 12 человек);

– существует большая вероятность получения случайного выбора из-за аморфной системы отношений выбирающего человека с окружающими или же из-за ложных ответов, с целью продемонстрировать нормальную лояльность к окружающим и к экспериментатору.

Этап IV включает в себя разработку **социометрической карточки (социометрической анкеты)**. Конечный результат разработки будет зависеть от того, (1) является ли социометрия самостоятельным исследованием, либо (2) представляет собою часть расширенного анкетного опроса.

В первом случае порядок составления карточки может быть таким:

а) готовятся списки членов группы; каждому члену группы присваивается порядковый номер, под которым он регистрируется в списке (этот номер становится шифром респондента);

б) готовится сама карточка, которая должна содержать следующие элементы:

- заголовок;
- инструкцию (в которой указывается допустимое количество необходимых выборов; способы заполнения граф выборов и т. п.);
- необходимые подписи.

Во втором случае, когда социометрическое исследование является частью более расширенного опроса, социометрические критерии могут быть включены в общий опросный лист отдельным подразделом.

Пример социометрической карточки представлен таблицей 5.1. В такой карточке каждый член группы указывает свое отношение к другим членам группы по ряду критериев (например, с точки зрения совместной работы, участия в решении деловой задачи, проведения досуга и т. д.) Критерии определяются в соответствии с программой исследования, его целями и задачами.

При опросе без ограничения выборов в социометрической карточке после каждого критерия должна быть выделена графа, размеры которой позволили бы давать достаточно полные ответы. При опросе с ограничением выборов справа от каждого критерия

Таблица 5.1

Пример социометрической карточки

Социометрическая карточка						
№	Тип	Критерии	Выборы			
1.	Работа	а) Кого бы вы хотели выбрать своим бригадиром? б) Кого бы вы не хотели выбрать своим бригадиром?				
2.	Досу	а) Кого бы вы хотели пригласить на встречу Нового года? б) Кого бы вы не хотели пригласить на встречу Нового года?				

на карточке чертится столько вертикальных граф, сколько выборов предполагается разрешить сделать в данной группе. Определение числа выборов для разных по численности групп, можно осуществить, используя специальную таблицу, представленную ниже (см. Табл. 5.2).

Таблица 5.2

Таблица величин, ограничивающих социометрические выборы

Число членов групп	Социометрическое ограничение (d)	Вероятность случайного выбора $P(A)$
5–7	1	0,20–0,14
8–11	2	0,25–0,18
12–16	3	0,25–0,19
17–21	4	0,23–0,19
22–26	5	0,22–0,19
27–31	6	0,22–0,19
32–36	7	0,21–0,19

Этап V представляет собою собственно проведение социометрии. В первую очередь, осуществляется инструктаж респондентов, а затем проводится сам опрос. Правильно проведенный инструктаж является залогом высокого качества полученной в результате исследования информации. Желательно, чтобы инструкция содержала мотивационную часть (освещение проблемы, на разрешение которой направлено исследование), чтобы каждый респондент осознавал значимость той информации, которую получает от него исследователь.

В целом сама по себе социометрическая процедура является достаточно серьезным эмоциональным испытанием, особенно для тех, кто занимает крайние статусные позиции: самые высокие (лидерские) и самые низкие (аутсайдерские). В этой связи человек, проводящий социометрический опрос, должен постараться снизить общую напряженность ситуации, установить доверительный контакт с группой, в доброжелательной, неформальной, спокойной форме ознакомить респондентов с основными положениями инструкции, «мягко» сформулировать цели и задачи.

Этап VI предполагает обработку данных и интерпретацию результатов социометрического исследования. Простейшими способами количественной обработки данных являются табличный (построение социометрической матрицы), индексологический (нахождение социометрических индексов) и графический (вычерчивание социограммы). Каждый из них с подробностями будет рассмотрен в трех последующих подразделах.

5.3. Изучение внутригрупповой структуры, выявление «социометрических позиций» членов группы. Построение социоматрицы и вычисление персональных социометрических индексов

При характеристике динамических процессов в малых группах, естественно, возникает вопрос о том, как группа организуется, кто берет на себя функции ее организации, каков психологический рисунок деятельности по управлению группой? Проблема лидерства и руководства является одной из кардинальных проблем социальной психологии, ибо оба эти процесса не просто относятся к проблеме интеграции групповой деятельности, а психологически описывают субъекта этой интеграции. Когда проблема обозначается как «проблема лидерства», то этим лишь отдается дань социально-психологической традиции, связанной с исследованием данного феномена. В современных условиях проблема должна быть поставлена значительно шире, как проблема руководства группой. Поэтому крайне важно

сделать, прежде всего, терминологические уточнения и развести понятия «лидер» и «руководитель». В украинском и русском языках для обозначения этих двух различных явлений существует два специальных термина (также, впрочем, как и в немецком, но не в английском языке, где «лидер» употребляется в обоих обозначенных случаях) и определены различия в содержании этих понятий. При этом не рассматривается употребление понятия «лидер» в политической терминологии. Б. Д. Парыгин называет следующие различия лидера и руководителя: 1) лидер, в основном, призван осуществлять регуляцию межличностных отношений в группе, в то время как руководитель осуществляет регуляцию официальных отношений группы как некоторой социальной организации; 2) лидерство можно констатировать в условиях микросреды (каковой и является малая группа), руководство – элемент макросреды, то есть оно связано со всей системой общественных отношений; 3) лидерство возникает стихийно, руководитель всякой реальной социальной группы либо назначается, либо избирается, но так или иначе этот процесс не является стихийным, а, напротив, целенаправленным, осуществляемым под контролем различных элементов социальной структуры; 4) явление лидерства менее стабильно, выдвижение лидера в большой степени зависит от настроения группы, в то время как руководство – явление более стабильное; 5) руководство подчиненными в отличие от лидерства обладает гораздо более определенной системой различных санкций, которых в руках лидера нет; 6) процесс принятия решения руководителем (и вообще в системе руководства) значительно более сложен и опосредован множеством различных обстоятельств и соображений, не обязательно коренящихся в данной группе, в то время как лидер принимает более непосредственные решения, касающиеся групповой деятельности; 7) сфера деятельности лидера – в основном малая группа, где он и является лидером, сфера действия руководителя шире, поскольку он представляет малую группу в более широкой социальной системе. Эти различия (с некоторыми вариантами) называют и другие авторы.

Как видно из приведенных соображений, лидер и руководи-

тель имеют, тем не менее, дело с однопорядковым типом проблем, а именно, они призваны стимулировать группу, нацеливать ее на решение определенных задач, заботиться о средствах, при помощи которых эти задачи могут быть решены. Хотя по происхождению лидер и руководитель различаются, в психологических характеристиках их деятельности существуют общие черты, что и дает право при рассмотрении проблемы зачастую описывать эту деятельность как идентичную, хотя это, строго говоря, не является вполне точным. Лидерство есть чисто психологическая характеристика поведения определенных членов группы, руководство в большей степени есть социальная характеристика отношений в группе, прежде всего с точки зрения распределения ролей управления и подчинения. В отличие от лидерства руководство выступает как регламентированный обществом правовой процесс. Чтобы изучить психологическое содержание деятельности руководителя, можно использовать знание механизма лидерства, но одно знание этого механизма ни в коем случае не дает полной характеристики деятельности руководителя. Поэтому последовательность в анализе данной проблемы должна быть именно такой: сначала выявление общих характеристик механизма лидерства, а затем интерпретация этого механизма в рамках конкретной деятельности руководителя.

Лидером является такой член малой группы, который выдвигается в результате взаимодействия членов группы для организации группы при решении конкретной задачи. Он демонстрирует более высокий, чем другие члены группы, уровень активности, участия, влияния в решении данной задачи. Таким образом, лидер выдвигается в конкретной ситуации, принимая на себя определенные функции. Остальные члены группы принимают лидерство, то есть строят с лидером такие отношения, которые предполагают, что он будет вести, а они будут ведомыми. Лидерство необходимо рассматривать как групповое явление: лидер немалозначим в одиночку, он всегда дан как элемент групповой структуры, а лидерство есть система отношений в этой структуре. Поэтому феномен лидерства относится к динамическим процессам малой группы. Этот процесс может быть достаточно противо-

речивым: мера притязаний лидера и мера готовности других членов группы принять его ведущую роль могут не совпадать.

Выяснить действительные возможности лидера – значит выяснить, как воспринимают лидера другие члены группы. Мера влияния лидера на группу также не является величиной постоянной, при определенных обстоятельствах лидерские возможности могут возрастать, а при других, напротив, снижаться. Иногда понятие лидера отождествляется с понятием «авторитет», что не вполне корректно: конечно, лидер выступает как авторитет для группы, но не всякий авторитет обязательно означает лидерские возможности его носителя. Лидер должен организовать решение какой-то задачи, авторитет такой функции не выполняет, он просто может выступать как пример, как идеал, но вовсе не брать на себя решение задачи. Поэтому феномен лидерства – это весьма специфическое явление, не описываемое никакими другими понятиями.

Определение лидеров и аутсайдеров в группе осуществляется путем нахождения персональных социометрических индексов (П.С.И.) с помощью нехитрых математических вычислений. Однако перед этим обычно строится социоматрица, как инструмент, облегчающий дальнейшую работу исследователя с данными.

Социоматрица. Социоматрица – это матрица связей, с помощью которой анализируются внутриколлективные отношения. В социоматрицу в форме числовых значений и символов заносится информация, полученная в ходе опроса. Анализ социоматрицы по каждому критерию дает достаточно наглядную картину взаимоотношений в группе. Могут быть построены суммарные социоматрицы, дающие картину выборов по нескольким критериям, а также социоматрицы по данным межгрупповых выборов.

Основное достоинство социоматрицы – возможность представить выборы в числовом виде, что, в свою очередь, позволяет проранжировать порядок влияний в группе. На основе социоматрицы строится социограмма – карта социометрических выборов (социометрическая карта), производится расчет социометрических индексов.

Результаты выборов разносятся по матрице с помощью условных обозначений. Таблицы результатов заполняются в первую очередь, в отдельности по деловым и личным отношениям. По вертикали записываются за номерами фамилии всех членов группы, которая изучается; по горизонтали – только их номер. На соответствующих пересечениях цифрами +1, +2, +3 обозначают тех, кого выбрал каждый испытуемый в первую, вторую, третью очередь, цифрами –1, –2, –3 – тех, кого подопытный не избирает в первую, вторую и третью очередь. Взаимный положительный или отрицательный выбор выделяется, каким-либо заранее установленным способом (независимо от очередности выбора). После того как положительные и отрицательные выборы будут занесены в таблицу, нужно подсчитать по вертикали алгебраическую сумму всех полученных каждым членом группы выборов (сумма выборов). Потом – подсчитать сумму баллов для каждого члена группы, учитывая при этом, что выбор в первую очередь равняется +3(–3) баллам, во вторую – +2(–2), в третью – +1(–1). После этого подсчитывается общая алгебраическая сумма, которая и определяет статус в группе.

Например, академик РАН Г. В. Осипов в одном из своих учебных пособий, посвященном описанию методов измерения в социологии, говорит о том, что социометрическая техника (как социометрическая мера) – это некоторый инструмент оценки предпочтения и неpreferенция индивидов в группе [3, с. 68–69]. Как простейший пример можно взять квадратную таблицу.

	A	B	C	D	E
A		1	1	0	0
B	0		0	1	0
C	0	1		0	1
D	1	0	1		0
E	1	1	0	1	

Как уже писалось, вдоль левого столбца и верхней строки данной матрицы записаны обозначения N индивидов группы. В клетке по строке пишем «1», если индивид i предпочитает индивида j , и «0», если не предпочитает. Клетки по диагонали можно заполнить единицами либо заштриховать (как в случае, представленном выше). Если мы суммируем столбцы, то получим величины, определяющие число предпочтений данного индивида. Можно считать их баллами индивидов по предпочтению и проранжировать индивидов по этим баллам.

В дальнейшем при использовании социометрической таблицы стали применяться модификации. Можно, например, ввести три показателя – предпочтение, нейтральность, не предпочтение и, соответственно, в клетках будет стоять одно из трех чисел – 1, 2, 3 (или 1, 0 – 1) и т. д.

С помощью социометрической таблицы можно ввести разнообразные индексы, характеризующие структуру всей группы. Таких индексов множество. Каждый исследователь вводил свой. Некоторые из них представляют интерес. Можно ввести отношение предпочтения как число всех предпочтений и не предпочтений, а также индекс совместимости группы как число взаимных предпочтений, деленное на объем группы.

Анализ социоматрицы по каждому критерию дает достаточно наглядную картину взаимоотношений в группе. Могут быть построены суммарные социоматрицы, дающие картину выборов по нескольким критериям, а также социоматрицы по данным межгрупповых выборов. Основное достоинство социоматрицы – возможность представить выборы в числовом виде, что в свою очередь позволяет проранжировать членов группы по числу полученных и предпочтенных выборов.

На основе данных социоматрицы могут быть осуществлены математические вычисления по нахождению персональных социометрических индексов для каждого члена группы («П.С.И.»). П.С.И. характеризуют его индивидуальные социально-психологические свойства и личностные качества. Основными П.С.И. считаются: индекс социометрического статуса i -члена; эмоциональной экспансивности j -члена, объема, интенсивности

и концентрации взаимодействия ij -члена. Символы i и j обозначают одно и то же лицо, но в разных ролях; i – выбираемый, j – он же выбирающий, ij – совмещение ролей.

Индекс социометрического статуса i -члена группы определяется по формуле:

$$C_i = \frac{\sum_{j=1}^N (R_i^+ + R_i^-)}{N - 1},$$

где • C_i – социометрический статус i -члена;

• R_i (+, –) – полученные i -членом выборы;

• N – число членов группы.

Возможен также расчет ***С-положительного и С-отрицательного статуса*** в группах малой численности (N).

Индекс эмоциональной экспансивности j -члена группы характеризует потребность личности в общении и высчитывается по формуле:

$$E_j = \frac{\sum_{i=1}^N (R_{ji}^+ + R_{ji}^-)}{N - 1},$$

где • E_j – эмоциональная экспансивность j -члена;

• R_j (+, –) – сделанные j -членом выборы;

• N – число членов группы.

5.4. Изучение степени сплоченности/разобщенности группы. Примеры вычисления групповых социометрических индексов

Возникнув благодаря внешним обстоятельствам, малая группа «переживает» длительный процесс своего становления в качестве социальной общности. С социологической точки зрения, важнейшим содержанием этого процесса является развитие групповой сплоченности. В ходе этого развития группа не просто продуцирует некоторые нормы и ценности, а члены

ее не просто усваивают их. Осуществляется гораздо более глубокая интеграция группы, когда ценности о предметной деятельности группы все в большей степени разделяются отдельными индивидами, не потому, что они им больше или меньше «нравятся», а потому, что индивиды включены в саму их совместную деятельность. Деятельность же эта становится столь значимой в жизни каждого члена группы, что он принимает ее ценности не под влиянием развития коммуникаций, убеждения, но самым фактом своего все более полного и активного включения в деятельность группы. Главной детерминантной образования группы в социально-психологическом значении этого слова выступает совместная деятельность индивидов как членов данной группы. Эта деятельность – не только внешне заданное условие существования данной группы, но и внутреннее основание ее существования [8].

Проблема групповой сплоченности имеет солидную традицию ее исследования, которая опирается на понимание группы, прежде всего, как некой системы межличностных отношений, имеющих эмоциональную основу. Несмотря на наличие разных вариантов интерпретации сплоченности, эта общая исходная посылка присутствует во всех случаях.

На практическом уровне все представления о малой группе как единой социальной общности складываются посредством формирования целостной конфигурации выборов в группе, путем вычисления общих для всей группы индексов. Так, в русле социометрического направления, сплоченность прямо связывалась с таким уровнем развития межличностных отношений, для которого характерен высокий процент выборов, основанных на взаимной симпатии. Социометрия предложила специальный «индекс групповой сплоченности», который вычисляется как отношение числа взаимных положительных выборов к общему числу возможных выборов [13, с. 90]:

$$C_{\text{гр}} = \frac{\sum r^{+}}{1/2 N(N-1)},$$

- где • C_{gr} – сплоченность;
- r^+ – положительный выбор;
 - N – число членов группы.

Содержательная характеристика взаимных положительных выборов здесь, как и вообще при применении социометрической методики, опущена. Индекс групповой сплоченности – строго формальная характеристика малой группы.

Другой подход был предложен Л. Фестингером, когда сплоченность анализировалась на основе частоты и прочности коммуникативных связей, обнаруживаемых в группе. Буквально сплоченность определялась как «сумма всех сил, действующих на членов группы, чтобы удерживать их в ней». Влияние школы К. Левина на Л. Фестингера придало особое содержание этому утверждению: «силы» интерпретировались либо как привлекательность группы для индивида, либо как удовлетворенность членством в группе. Но и привлекательность, и удовлетворенность анализировались при помощи выявления чисто эмоционального плана отношений в группе, поэтому, несмотря на иной по сравнению с социометрией подход, сплоченность и здесь представлялась как некоторая характеристика системы эмоциональных предпочтений членов группы.

Была предложена еще одна программа исследования сплоченности, связанная с работами Т. Ньюкома, который ввел особое понятие «согласия», с помощью которого и предложил интерпретировать сплоченность. Ученый выдвигает идею, не похожую на те, которые содержались в подходах Я. Морено и Л. Фестингера, а именно идею необходимости возникновения сходных ориентаций членов группы, по отношению к тем либо иным значимым для них ценностям. Несомненная продуктивность этой идеи, к сожалению, оказалась девальвированной, поскольку дальнейшее ее развитие попало в жесткую схему теории поля. Развитие сходных ориентаций, то есть достижение согласия, мыслилось как снятие напряжений в поле действия индивидов, причем это «снятие» происходило на основе определенных эмоциональных реакций индивидов. Хотя и с оговорками, но

мысль об эмоциональной основе сплоченности оказалась основополагающей и в этом варианте объяснения.

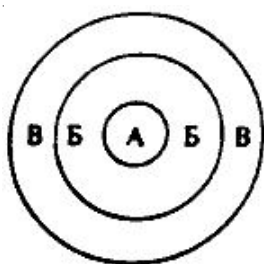
Существует целый ряд экспериментальных работ по выявлению групповой сплоченности (группового единства). Из них надо назвать исследования А. Бейвеласа, в которых особое значение придается характеру групповых целей. Операциональные цели группы – это цели, которые «работают» на построение оптимальной системы коммуникаций; символические цели группы – это цели, соответствующие индивидуальным намерениям членов группы. Сплоченность зависит от реализации и того, и другого характера целей. Как видим, интерпретация феномена становится здесь богаче.

Логично представить себе новый подход к исследованию сплоченности, если он будет опираться на принятые принципы понимания группы и, в частности, на идею о том, что главным интегратором группы является совместная деятельность ее членов. Тогда процесс формирования группы и ее дальнейшего развития предстает как процесс все большего сплачивания этой группы, но отнюдь не на основе увеличения лишь эмоциональной ее привлекательности, а на основе все большего включения индивидов в процесс совместной деятельности. Для этого выявляются иные основания сплоченности. Чтобы лучше понять их природу, следует сказать, что речь идет именно о сплоченности группы, а не о совместимости людей в группе. Хотя совместимость и сплоченность тесно связаны, каждое из этих понятий обозначает разный аспект характеристики группы. Совместимость членов группы означает, что данный состав группы возможен для обеспечения выполнения группой ее функций, что члены группы могут взаимодействовать. Сплоченность группы означает, что данный состав группы не просто возможен, но что он интегрирован наилучшим образом, что в нем достигнута особая степень развития отношений, а именно такая степень, при которой все члены группы в наибольшей мере разделяют цели групповой деятельности и те ценности, которые связаны с этой деятельностью. Это отличие сплоченности от совместимости подвело нас к пониманию существа сплоченности в рамках принципа деятельности.

В социально-психологической мысли советских лет новые принципы исследования сплоченности разрабатывались российским психологом А. В. Петровским. На сегодняшний день они составляют часть единой концепции, названной ранее «стратометрической концепцией групповой активности», а позднее — «теорией деятельностного опосредования межличностных отношений в группе». Основная идея заключается в том, что всю структуру малой группы можно представить себе как состоящую из трех (в последней редакции четырех) основных слоев, или, в иной терминологии, «страт»: внешний уровень групповой структуры, где даны непосредственные эмоциональные межличностные отношения, то есть то, что традиционно измерялось социометрией; второй слой, представляющий собой более глубокое образование, обозначаемое термином «ценностно-ориентационное единство» (ЦОЕ), которое характеризуется тем, что отношения здесь опосредованы совместной деятельностью, выражением чего является совпадение для членов группы ориентации на основные ценности, касающиеся процесса совместной деятельности. Социометрия, построив свою методику на основе выбора, не показывала, как отмечалось, мотивов этого выбора. Для изучения второго слоя (ЦОЕ) нужна иная методика, позволяющая вскрыть мотивы выбора. Теория же дает ключ, при помощи которого эти мотивы могут быть обнаружены: это — совпадение ценностных ориентаций, касающихся совместной деятельности. Третий слой групповой структуры расположен еще глубже и предполагает еще большее включение индивида в совместную групповую деятельность: на этом уровне члены группы разделяют цели групповой деятельности, и, следовательно, здесь могут быть выявлены наиболее серьезные, значимые мотивы выбора членами группы друг друга. Можно предположить, что мотивы выбора на этом уровне связаны с принятием общих ценностей, но более абстрактного уровня: ценностей, связанных с более общим отношением к труду, к окружающим, к миру. Этот третий слой отношений был назван «ядром» групповой структуры.

Все сказанное имеет непосредственное отношение к понима-

нию сплоченности группы как определенного процесса развития внутригрупповых связей, который соответствует развитию групповой деятельности. Три слоя групповых структур могут одновременно быть рассмотрены и как три уровня развития группы, в частности, три уровня развития групповой сплоченности (см. *Рис. 5.1*). На первом уровне (что соответствует поверхностному слою внутригрупповых отношений) сплоченность действительно выражается развитием эмоциональных контактов (**В**). На втором уровне (что соответствует второму слою – ЦОЕ) происходит дальнейшее сплочение группы, и теперь это выражается в совпадении у членов группы основной системы ценностей, связанных с процессом совместной деятельности (**Б**). На третьем уровне (что соответствует «ядерному» слою внутригрупповых отношений) интеграция группы (а значит, и ее сплоченность) проявляется в том, что все члены группы начинают разделять общие цели групповой деятельности (**А**).



слой В – непосредственные эмоциональные контакты;
слой Б – ЦОЕ (ценностно-ориентационное единство);
слой А – «ядро» (совместная групповая деятельность и ее цели)

Рис. 5.1. Структура малой социальной группы, по уровню сплоченности (с точки зрения стратометрической концепции)

Существенным моментом при этом выступает то обстоятельство, что развитие сплоченности осуществляется не за счет развития лишь коммуникативной практики (как у Ньюкома), но на основе совместной деятельности. Кроме того, единство группы, выраженное в единстве ценностных ориентаций членов группы, интерпретируется не просто как сходство этих ориентаций, но и как воплощение этого сходства в ткань практических действий членов группы. При такой интерпретации сплоченности обязателен третий шаг в анализе, то есть переход от установления

единства ценностных ориентации к установлению еще более высокого уровня единства – единства целей групповой деятельности как выражения сплоченности. Как пишет известный российский ученый, доктор психологических наук А. И. Донцов²²: «...Если каждый из вышеназванных феноменов сплоченности является показателем интегрированности лишь отдельных пластов и слоев внутригрупповой активности, то общность цели, будучи детерминантой всех их вместе взятых, может служить референтом действительного единства группы как целого...» [11; 12]. Можно считать, конечно, что совпадение целей групповой деятельности есть в то же самое время и высший уровень ценностного единства группы, поскольку сами цели совместной деятельности есть также определенная ценность. Таким образом, в практике исследования сплоченность должна быть проанализирована и как совпадение ценностей, касающихся предмета совместной деятельности, и как своего рода «деятельностное воплощение» этого совпадения.

Эксперимент А. И. Донцова, проведенный в 14 московских средних школах, имел целью выявить, как в работе учителей совпадают представления о ценности деятельности с реальным воплощением их в повседневной практике преподавания и воспитания. Выявлялись два пласта сплоченности: сплоченность, демонстрируемая при оценке «эталонного» ученика, и сплоченность, демонстрируемая при оценке реальных учеников. В результате исследования был получен вывод, что общность оценок реальных учеников, данная учителями одной школы, выше, чем согласованность их представлений об эталоне ученика, то есть сплоченность в реальной деятельности оказалась выше, чем сплоченность, регистрируемая лишь как совпадение мнений (ибо отношением к эталону может быть только мнение, но не реальная деятельность).

²² Донцов А.И. (род. 15 октября 1949 г.) – российский психолог, д-р психологических наук, профессор, действительный член Российской академии образования.

Вторая часть исследования дала довольно любопытный результат. Когда оцениванию подвергались не ученики, а коллеги-учителя, то единство ценностных представлений оказалось выше в том случае, когда речь шла именно об эталоне учителя, и ниже, когда давались оценки реальным коллегам. Интерпретация этого факта снова подтверждает основной принцип: конкретным предметом деятельности учителя не является другой учитель – коллега, поэтому оценивание его, а значит, и совпадение такого рода ценностей не есть параметр непосредственной конкретной деятельности данной группы. Напротив, эталонный коллега в большей степени оказался ценностью, включенной в непосредственную практику работы учителя. (Образ такого «эталонного» коллеги возникает, например, на различных методических конференциях, собраниях «предметников».) Таким образом, была подтверждена гипотеза исследования о том, что действительная интеграция группы (а, следовательно, и ее сплоченность) осуществляется, прежде всего, в ходе совместной деятельности [12].

Информацию о сплоченности группы, полученную с помощью вычисления $C_{гр.}$ и характеризующую группу как целостность, могут дополнить также другие Г.С.И., такие как индекс эмоциональной экспансивности группы и индекс психологической взаимности.

Индекс *эмоциональной экспансивности группы* вычисляется по формуле:

$$Ag = \frac{\sum_{i=1}^N (\sum_{j=1}^N R_j^{(+,-)})}{N},$$

где • Ag – экспансивность группы;

• N – число членов группы;

• $R_j (+, -)$ – сделанные j -членом выборы.

Индекс показывает среднюю активность группы при решении задачи социометрического теста (в расчете на каждого члена группы).

Индекс психологической взаимности («сплоченности группы») в группе высчитывается по формуле:

$$G_g = \frac{\sum_{ij=1}^N (\sum_{jj=1}^N A_{ij}^+)}{1/2 \cdot N(N-1)},$$

где • G_g – взаимность в группе по результатам положительных выборов;

- A_{ij}^+ – число положительных взаимных связей в группе;
- N – число членов группы.

Подводя итог сказанному, подчеркнем, что составление социоматриц, а также вычисление персональных и групповых социометрических индексов являются простейшими способами количественной обработки данных социометрического опроса. Однако, как правило, исследователь не ограничивается только лишь таблицами и индексами. Более глубокий анализ взаимоотношений в группе предполагает обнаружение внутригрупповых подсистем, что требует дополнительных процедур по работе с данными и привлечения иных математических методов. Об одном из таких методов, заключающемся в построении графов, на основе информации, представленной в социоматрице и зафиксированной в индексах, и пойдет речь в следующем подразделе.

5.5. Обнаружение внутригрупповых подсистем по результатам социометрического исследования. Построение социограмм с обращением к теории графов

Когда социометрические карточки заполнены и собраны, начинается этап их математической обработки. Простейшими способами количественной обработки данных социометрического опроса являются табличный и индексологический методы, описанные с подробностями в двух предыдущих параграфах.

Существенное дополнение к этим методам представляет социограммная техника. Она дает возможность более глубокого качественного описания и наглядного представления групповых явлений, позволяет обнаружить внутригрупповые подсистемы, существующие в группе.

Социограмма – графическое изображение реакции респондентов (участников социометрического опроса) друг на друга при ответах на социометрический критерий. Она позволяет произвести сравнительный анализ структуры взаимоотношений в группе в пространстве на некоторой плоскости («щите») с помощью специальных знаков (рис. ниже). Она даёт наглядное представление о внутригрупповой дифференциации членов группы за их статусом (популярностью).

Построение социограммы начинается с поиска центральных, наиболее влиятельных членов, а затем выявления взаимных пар и группировок. Группировки состояются из взаимосвязанных лиц, стремящихся выбирать друг друга. При построении социограммы исследователи довольно часто апеллируют к теории графов. С помощью этой теории обычно исследуются свойства конечных множеств с заданными отношениями между их элементами. Так как численность малой социальной группы – конечное множество и отношения между членами этой группы преимущественно определены (с помощью табличного и индексологического методов), вполне обоснованным представляется обращение к основным положениям теории графов для осуществления более глубокого анализа взаимоотношений в группе, их дифференциации и иерархизации.

Как известно теория графов – раздел дискретной математики, изучающий свойства графов²³. Формальное, так сказать, «математическое» определение графа таково: задано конечное множество X , состоящее из n элементов ($X = \{1, 2, \dots, n\}$), называемых вершинами графа, и подмножество V декартова

²³ Начало теории графов датируют 1736 г., когда **Л. Эйлер** решил популярную в то время «задачу о кенигсбергских мостах». Термин «граф» впервые был введен спустя 200 лет (в 1936 г.) **Д. Кенигом**.

произведения $X' \times X$, называемое множеством дуг, тогда ориентированным *графом* G называется совокупность (X, V) . Неориентированным графом называется совокупность множества X и множества неупорядоченных пар элементов, каждый из которых принадлежит множеству X . Дугу между вершинами i и j будем обозначать (i, j) . Число дуг графа будем обозначать m ($V = (v_1, v_2, \dots, v_m)$).

Несмотря на такое, на первый взгляд, сложное и специфическое определение графа, язык графов универсален, он оказывается удобным для описания многих физических, технических, экономических, биологических, социальных и других систем, а следовательно, часто используется многими научными дисциплинами, не связанными напрямую с математикой. На междисциплинарном уровне граф может быть рассмотрен как система, состоящая из множества кружков и множества соединяющих эти кружки линий (геометрический способ задания графа). Кружки – это *вершины* графа, линии со стрелками – *дуги*, без стрелок – *ребра*. Граф, в котором направление линий не выделяется (все линии являются ребрами), называется *неориентированным*; граф, в котором направление линий принципиально (линии являются дугами) называется *ориентированным* [1; 9].

Таким образом, следуя основным положениям теории графов, каждого индивида-респондента можно обозначить кружком (точкой). Если он предпочитает/не предпочитает некоторого другого члена группы, то следует провести от этого кружка (точки) стрелку к кружку (точке), обозначающей другого индивида. Очевидно, что социограмма сразу показывает наглядно, какие индивиды являются наиболее предпочитаемыми лидерами, а какие – наименее предпочитаемыми аутсайдерами.

При этом условные обозначения предпочтений/отвержений могут иметь следующий вид:

- $\longrightarrow$ позитивный односторонний выбор;
- $\longleftrightarrow$ позитивный обоюдный выбор;
- $\dashrightarrow$ негативный односторонний выбор;
- $\dashleftarrow$ негативный обоюдный выбор.

С помощью этих обозначений, открываются возможности осуществления сравнительного анализа структуры взаимоотношений в группе на некоторой плоскости, так называемом «щите» (см. Рис. 5.2).



Рис. 5.2. Условные обозначения связей между членами группы на «щите»

Наиболее часто в социометрических измерениях встречаются положительные группировки из двух-трех членов, реже из четырех и более членов (см. Рис. 5.3).

Социограммы подразделяются на групповые и индивидуальные. Первые изображают картину взаимоотношений в группе в целом, вторые – систему отношений, существующих у интересующего исследователя индивида с остальными членами его группы.

Групповые социограммы: конвенциональная социограмма и социограмма-мишень. На *конвенциональной социограмме*

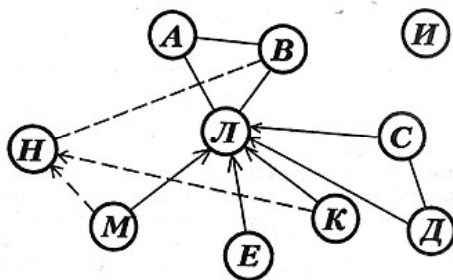


Рис. 5.3. Пример социогаммы, иллюстрирующий выделение группировок

индивиды, составляющие группу, изображаются в виде кружочков, соединенных между собой стрелками, символизирующими социометрические выборы или отклонения. При построении конвенциональной социогаммы индивиды располагаются по вертикали в соответствии с количеством полученных ими выборов таким образом, чтобы в верхней части социогаммы оказались те, кто получил наибольшее количество выборов. Индивидов необходимо располагать на таком расстоянии друг от друга, чтобы оно было пропорционально порядку выбора. Если, например, два индивида, А и Г, выбрали друг друга в первую очередь, то расстояние между изображающими их кружочками на рисунке должно быть минимальным; если индивид Д выбрал А в третью очередь, то длина стрелки, соединяющей А и Д, должна быть примерно в три раза больше, чем длина стрелки, соединяющей А и Г (см. Рис. 5.4).

Социогамма-мишень представляет собой систему концентрических окружностей, количество которых равно максимальному количеству выборов, полученных в группе. Все члены группы располагаются на окружностях, в соответствии с количеством полученных выборов. Социогамма-мишень делится на секторы по социально-демографическим характеристикам группы, например, по полу, возрасту и т. п. (см. Рис. 5.5).

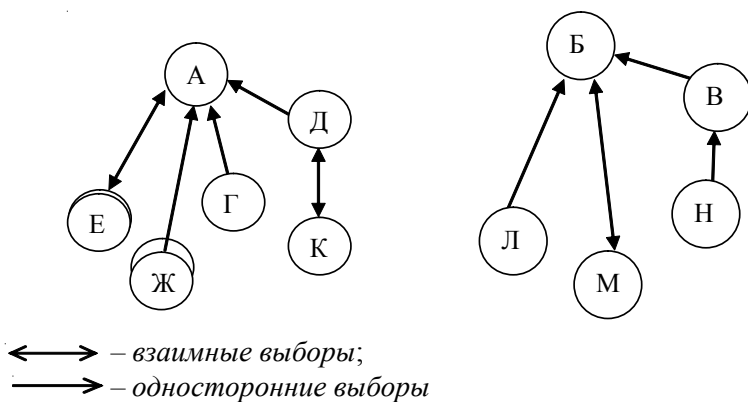


Рис. 5.4. Конвенциональная социограмма, изображающая отношения в группе из 11 человек

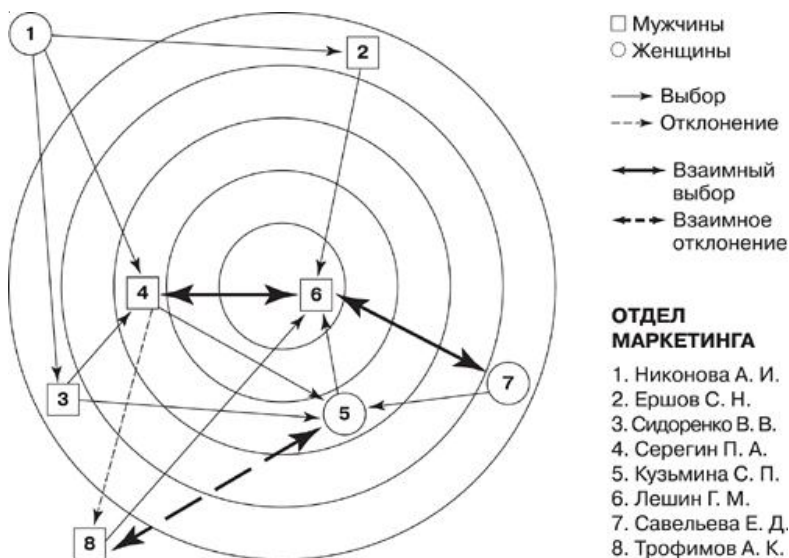


Рис. 5.5. Социограмма-мишень для рабочего коллектива из восьми человек

Изолированные

Игнорируемые

Предпочитаемые

Звезда

К

Б

А

Д

Ж

Е

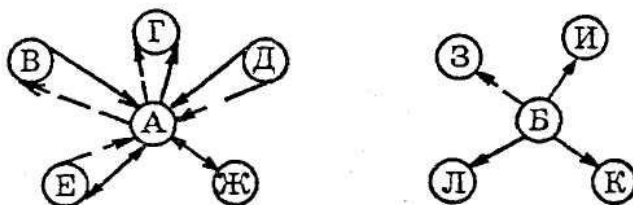
Г

Л

двухдольные социограммы. Для более

ческий опрос является интересным и довольно популярным методом изучения социального взаимодействия на микроуровне –

284



Индивидуальные социогаммы лидера (А)
и изолированного (Б):

- ↔ - взаимные выборы;
- - односторонние выборы;
- ↔ - ожидаемые взаимные выборы;
- - односторонние ожидаемые выборы

Рис. 5.7. Индивидуальные социогаммы
для лидера (А) и изолированного (Б)

межличностном, внутригрупповом (в малых социальных группах). Сама процедура проведения социометрии не требует больших временных затрат и, как правило, занимает не более 15 минут. Социометрия весьма полезна в прикладных исследованиях, особенно в работах по совершенствованию отношений в коллективе. Однако она не является радикальным способом разрешения внутригрупповых проблем, причины которых следует искать не в симпатиях и антипатиях членов группы, а в более глубоких источниках.

Вопросы для самоконтроля

1. Почему социологу может быть интересным изучение малых социальных групп и межличностных отношений?
2. Какие еще науки о человеке и обществе изучают малые социальные группы и межличностные отношения? В чем специфика социологического подхода к изучению?
3. Опишите, с помощью каких методов социологи могут исследовать малые социальные группы?
4. Дайте определения: социометрия, социоматрица, социогамма, социометрический индекс.

5. Каково предназначение и содержание социометрических методов, в чем заключается специфика социометрических опросов?

6. Объясните необходимость применения социометрического метода для изучения социальных процессов и явлений.

7. Охарактеризуйте сущность, разновидности и эвристические возможности социограммы и социоматрицы.

8. Покажите динамику внутригрупповых предпочтений, когда в группу попадает новый лидер.

9. Обоснуйте специфику взаимоотношений объект-объект в коэффициентах социометрического анализа.

10. Каковы основные социометрические индексы? Объясните их значения.

11. Каковы возможности и ограничения социометрического метода?

12. Подумайте и опишите случаи использования теории графов при исследовании конкретных социальных процессов или явлений, в разрешении тех или иных социологических проблем.

13. Перечислите и определите основные термины теории графов.

14. Какова история развития теории графов?

15. В каких случаях социологи обращаются к теории графов?

16. В чем различия между стрелочными и вершинными графами?

17. В чем заключается метод прогностного графа и как он может быть применим в эмпирической социологии?

ПРИЛОЖЕНИЯ

ПРИЛОЖЕНИЕ 1

ПРИМЕРЫ ВОПРОСОВ, СФОРМУЛИРОВАННЫХ В РАЗЛИЧНЫХ ШКАЛАХ ИЗМЕРЕНИЙ

А. Шкала наименований

1. Пожалуйста, укажите ваш пол:

- 1) мужской,
- 2) женский.

2. Выберете марки электронной продукции, которые вы обычно покупаете:

- 1) «Сони»,
- 2) «Панасоник»,
- 3) «Филипс»,
- 4) «Орион»,
- 5) «Другое».

Б. Шкала порядка

1. Пожалуйста, проранжируйте фирмы, производящие электронную продукцию, в соответствии с системой вашего предпочтения. Поставьте «1» фирме, которая занимает первое место в системе ваших предпочтений; «2» – второй и т. д.:

- 1) «Сони»,
- 2) «Панасоник»,
- 3) «Филипс»,
- 4) «Орион»,
- 5) «Другое».

2. Из каждой пары магазинов, занимающихся продажей электротехники, обведите кружком тот, который вы предпочитаете:

- 1) (1) «Метро» и (2) «Комфи»,
- 2) (2) «Комфи» и (3) «Электроленд»,
- 3) (3) «Электроленд» и (4) «Мой».

3. Что вы скажете о ценах в ТЦ «Метро»:

- 1) они выше, чем в ТЦ «Восторг»,
- 2) те же самые, что и в ТЦ «Восторг»,
- 3) ниже, чем в ТЦ «Восторг».

В. Шкала интервалов

1. Пожалуйста, оцените каждую марку товара с точки зрения его качества:

Марка	Рейтинг (обведите одну из цифр)	
	Очень низкое	Очень высокое
«Монблан»	1 2 3 4 5 6 7 8 9 10	
«Паркер»	1 2 3 4 5 6 7 8 9 10	
«Кросс»	1 2 3 4 5 6 7 8 9 10	

2. Укажите степень вашего согласия со следующими заявлениями, обведя одну из цифр:

Заявление	Совершенно не согласен	Полностью согласен
а. Я всегда стремлюсь делать выгодные покупки	1 2 3 4 5	
б. Я люблю проводить время вне дома	1 2 3 4 5	
в. Я люблю готовить	1 2 3 4 5	

Г. Шкала отношений

1. Пожалуйста, укажите ваш возраст _____ полных лет

2. Укажите, сколько приблизительно раз за последний месяц вы делали покупки в круглосуточном магазине в интервале времени от 22 до 24 часов:

0 1 2 3 4 5 другое число раз _____

3. Какова вероятность того, что при составлении завещания вы прибегнете к помощи юриста?

_____ процентов

ПРИЛОЖЕНИЕ 2

СТАТИСТИЧЕСКИЕ ТАБЛИЦЫ

Таблица А

Таблица значений критических точек стандартного нормального распределения для различных уровней значимости

α	0,01	0,025	0,05	0,10	0,20	0,30
$Z_{кр}$	2,3263	1,9600	1,6449	1,2816	0,8416	0,5244

Таблица Б

Критическое значение для χ^2 распределения

df	0.30	0.25	0.10	0.05	0.02	0.01
1	1,074	1,323	2,706	3,841	5,412	6,635
2	2,408	2,773	4,605	5,991	7,824	9,210
3	3,665	4,108	6,251	7,815	9,837	11,345
4	4,878	5,385	7,779	9,488	11,668	13,277
5	6,064	6,626	9,236	11,070	13,388	15,086
6	7,231	7,841	10,645	12,592	15,033	16,812
7	8,383	9,037	12,017	14,067	16,622	18,475
8	9,524	10,219	13,362	15,507	18,168	20,090
9	10,656	11,389	14,684	16,919	19,679	21,666
10	11,781	12,549	15,987	18,307	21,161	23,209
11	12,899	13,701	17,275	19,675	22,618	24,725
12	14,011	14,845	18,549	21,026	24,054	26,217
13	15,119	15,984	19,812	22,362	25,471	27,688
14	16,222	17,117	21,064	23,685	26,873	29,141
15	17,322	18,245	22,307	24,996	28,259	30,578
16	18,418	19,369	23,542	26,296	29,633	32,000
17	19,511	20,489	24,769	27,587	30,995	33,409
18	20,601	21,605	25,989	28,869	32,346	34,805
19	21,689	22,718	27,204	30,144	33,687	36,191
20	22,775	23,828	28,412	31,410	35,020	37,566
21	23,858	24,935	29,615	32,671	36,343	38,932
22	24,939	26,039	30,813	33,924	37,659	40,289
23	26,018	27,141	32,007	35,172	38,968	41,638
24	27,096	28,241	33,196	36,415	40,270	42,980
25	28,172	29,339	34,382	37,652	41,566	44,314
26	29,246	30,435	35,563	38,885	42,856	45,642
27	30,319	31,528	36,741	40,113	44,140	46,963
28	31,391	32,620	37,916	41,337	45,419	48,278
29	32,461	33,711	39,087	42,557	46,693	49,588
30	33,530	34,800	40,256	43,773	47,962	50,892

Таблица В

Критические значения для *t*-распределения Стьюдента

<i>Уровень значимости для одностороннего критерия</i>					
	<i>0,10</i>	<i>0,05</i>	<i>0,25</i>	<i>0,01</i>	<i>0,05</i>
<i>Уровень значимости для двустороннего критерия</i>					
	<i>0,20</i>	<i>0,10</i>	<i>0,05</i>	<i>0,02</i>	<i>0,01</i>
1	3,078	6,314	12,706	31,821	63,656
2	1,886	2,920	4,303	6,965	9,925
3	1,638	2,353	3,182	4,541	5,841
4	1,533	2,132	2,776	3,747	4,604
5	1,476	2,015	2,571	3,365	4,032
6	1,440	1,943	2,447	3,143	3,707
7	1,415	1,895	2,365	2,998	3,499
8	1,397	1,860	2,306	2,896	3,355
9	1,383	1,833	2,262	2,821	3,250
10	1,372	1,812	2,228	2,764	3,169
11	1,363	1,796	2,201	2,718	3,106
12	1,356	1,782	2,179	2,681	3,055
13	1,350	1,771	2,160	2,650	3,012
14	1,345	1,761	2,145	2,624	2,977
15	1,341	1,753	2,131	2,602	2,947
16	1,337	1,746	2,120	2,583	2,921
17	1,333	1,740	2,110	2,567	2,898
18	1,330	1,734	2,101	2,552	2,878
19	1,328	1,729	2,093	2,539	2,861
20	1,325	1,725	2,086	2,528	2,845
21	1,323	1,721	2,080	2,518	2,831
22	1,321	1,717	2,074	2,508	2,819
23	1,319	1,714	2,069	2,500	2,807
24	1,318	1,711	2,064	2,492	2,797
25	1,316	1,708	2,060	2,485	2,787
26	1,315	1,706	2,056	2,479	2,779
27	1,314	1,703	2,052	2,473	2,771
28	1,313	1,701	2,048	2,467	2,763
29	1,311	1,699	2,045	2,462	2,756
30	1,310	1,697	2,042	2,457	2,750
40	1,303	1,684	2,021	2,423	2,704
60	1,296	1,671	2,000	2,390	2,660
120	1,289	1,658	1,980	2,358	2,617
∞	1,282	1,645	1,960	2,327	2,577

**Критическое значение коэффициента ранговой корреляции
Спирмена r_s для двухсторонних и односторонних
критериев α (2) и α (1) соответственно**

α (2) α (1)	0,20 0,10	0,10 0,05	0,05 0,025	0,02 0,01	0,01 0,005
4	1,000	1,000			
5	0,800	0,900	1,000	1,000	
6	0,657	0,829	0,886	0,943	1,000
7	0,571	0,714	0,786	0,893	0,929
8	0,524	0,643	0,738	0,833	0,881
9	0,483	0,600	0,700	0,783	0,833
10	0,455	0,564	0,648	0,745	0,794
11	0,427	0,536	0,618	0,709	0,755
12	0,406	0,503	0,587	0,671	0,727
13	0,385	0,484	0,560	0,648	0,703
14	0,367	0,464	0,538	0,622	0,675
15	0,354	0,443	0,521	0,604	0,654
16	0,341	0,429	0,503	0,582	0,635
17	0,328	0,414	0,485	0,566	0,615
18	0,317	0,401	0,472	0,550	0,600
19	0,309	0,391	0,460	0,535	0,584
20	0,299	0,380	0,447	0,520	0,570
21	0,292	0,370	0,435	0,508	0,556
22	0,284	0,361	0,425	0,496	0,544
23	0,278	0,353	0,415	0,486	0,532
24	0,271	0,344	0,406	0,475	0,521
25	0,265	0,337	0,398	0,466	0,511
26	0,259	0,331	0,390	0,457	0,501
27	0,255	0,324	0,382	0,448	0,491
28	0,250	0,317	0,375	0,440	0,483
29	0,245	0,312	0,368	0,433	0,475
30	0,240	0,306	0,362	0,425	0,467
31	0,236	0,301	0,356	0,418	0,450
32	0,232	0,296	0,350	0,412	0,452
33	0,229	0,291	0,345	0,405	0,446
34	0,225	0,287	0,340	0,399	0,439
35	0,222	0,283	0,335	0,394	0,433
36	0,219	0,279	0,330	0,388	0,427
37	0,216	0,275	0,325	0,383	0,421
38	0,212	0,271	0,321	0,378	0,415
39	0,210	0,267	0,317	0,373	0,410

Продолжение таблицы Г

α (2) α (1)	0,20 0,10	0,10 0,05	0,05 0,025	0,02 0,01	0,01 0,005
40	0,207	0,264	0,313	0,368	0,405
41	0,204	0,261	0,309	0,364	0,400
42	0,202	0,257	0,305	0,359	0,395
43	0,199	0,254	0,301	0,355	0,391
44	0,197	0,251	0,298	0,351	0,386
45	0,194	0,248	0,294	0,347	0,382
46	0,192	0,246	0,291	0,343	0,377
47	0,190	0,243	0,283	0,340	0,374
48	0,188	0,240	0,285	0,336	0,370
49	0,186	0,238	0,282	0,333	0,368
50	0,184	0,235	0,279	0,329	0,363
52	0,180	0,231	0,274	0,323	0,356
54	0,177	0,226	0,268	0,317	0,349
56	0,174	0,222	0,264	0,311	0,343
58	0,171	0,218	0,259	0,306	0,337
60	0,168	0,214	0,255	0,300	0,331
62	0,165	0,211	0,250	0,296	0,326
64	0,162	0,207	0,246	0,291	0,321
66	0,160	0,204	0,243	0,287	0,316
68	0,157	0,201	0,239	0,282	0,311
70	0,155	0,198	0,235	0,278	0,307
72	0,153	0,195	0,232	0,274	0,303
74	0,151	0,193	0,229	0,271	0,299
76	0,149	0,190	0,226	0,267	0,295
78	0,147	0,188	0,223	0,264	0,291
80	0,145	0,185	0,220	0,260	0,287
82	0,143	0,183	0,217	0,257	0,284
84	0,141	0,181	0,215	0,254	0,280
86	0,139	0,179	0,212	0,251	0,277
88	0,138	0,176	0,210	0,248	0,274
90	0,136	0,174	0,207	0,245	0,271
92	0,135	0,173	0,205	0,243	0,268
94	0,133	0,171	0,203	0,240	0,265
96	0,132	0,169	0,201	0,238	0,262
98	0,130	0,167	0,199	0,235	0,260
100	0,129	0,165	0,197	0,233	0,257

Преобразование коэффициента корреляции $z = \operatorname{arctg}(r)$

r	z	r	z	r	z
0	0,0000	0,45	0,4847	0,9	1,4722
0,01	0,0100	0,46	0,4973	0,91	1,5275
0,02	0,0200	0,47	0,5101	0,92	1,5890
0,03	0,0300	0,48	0,5230	0,93	1,6584
0,04	0,0400	0,49	0,5361	0,94	1,7380
0,05	0,0500	0,5	0,5493	0,95	1,8318
0,06	0,0601	0,51	0,5627	0,96	1,9459
0,07	0,0701	0,52	0,5763	0,961	1,9588
0,08	0,0802	0,53	0,5901	0,962	1,9721
0,09	0,0902	0,54	0,6042	0,963	1,9857
0,1	0,1003	0,55	0,6184	0,964	1,9996
0,11	0,1104	0,56	0,6328	0,965	2,0139
0,12	0,1206	0,57	0,6475	0,966	2,0287
0,13	0,1307	0,58	0,6625	0,967	2,0439
0,14	0,1409	0,59	0,6777	0,968	2,0595
0,15	0,1511	0,6	0,6931	0,969	2,0756
0,16	0,1614	0,61	0,7089	0,97	2,0923
0,17	0,1717	0,62	0,7250	0,971	2,1095
0,18	0,1820	0,63	0,7414	0,972	2,1273
0,19	0,1923	0,64	0,7582	0,973	2,1457
0,2	0,2027	0,65	0,7753	0,974	2,1649
0,21	0,2132	0,66	0,7928	0,975	2,1847
0,22	0,2237	0,67	0,8107	0,976	2,2054
0,23	0,2342	0,68	0,8291	0,977	2,2269
0,24	0,2448	0,69	0,8480	0,978	2,2494
0,25	0,2554	0,7	0,8673	0,979	2,2729
0,26	0,2661	0,71	0,8872	0,98	2,2976
0,27	0,2769	0,72	0,9076	0,981	2,3235
0,28	0,2877	0,73	0,9287	0,982	2,3507
0,29	0,2986	0,74	0,9505	0,983	2,3796
0,3	0,3095	0,75	0,9730	0,984	2,4101
0,31	0,3205	0,76	0,9962	0,985	2,4427
0,32	0,3316	0,77	1,0203	0,986	2,4774
0,33	0,3428	0,78	1,0454	0,987	2,5147
0,34	0,3541	0,79	1,0714	0,988	2,5550
0,35	0,3654	0,8	1,0986	0,989	2,5987
0,36	0,3769	0,81	1,1270	0,99	2,6467
0,37	0,3884	0,82	1,1568	0,991	2,6996
0,38	0,4001	0,83	1,1881	0,992	2,7587
0,39	0,4118	0,84	1,2212	0,993	2,8257
0,4	0,4236	0,85	1,2562	0,994	2,9031
0,41	0,4356	0,86	1,2933	0,995	2,9915
0,42	0,4477	0,87	1,3331	0,996	3,1063
0,43	0,4599	0,88	1,3758	0,997	3,2504
0,44	0,4722	0,89	1,4219	0,998	3,4534

Таблица E

Таблица значений интеграла вероятностей $\Phi(x) = \frac{1}{\sqrt{2\pi}} \int_0^x e^{-\frac{t^2}{2}} dt$

x	$\Phi(x)$	x	$\Phi(x)$	x	$\Phi(x)$	x	$\Phi(x)$
0,00	0,0000	1,00	0,3413	2,00	0,4772	3,00	0,4986
01	0040	01	3438	01	4778	01	4987
02	0080	02	3461	02	4783	02	4987
03	0120	03	3485	03	4788	03	4988
04	0160	04	3508	04	4793	04	4988
05	0199	05	3531	05	4798	05	4989
06	0239	06	3554	06	4803	06	4989
07	0279	07	3577	07	4808	07	4989
08	0319	08	3599	08	4812	08	4990
09	0359	09	3621	09	4817	09	4990
0,10	0,0398	1,10	0,3643	2,10	0,4821	3,10	0,4990
11	0438	11	3665	11	4825	11	4991
12	0478	12	3686	12	4830	12	4991
13	0517	13	3708	13	4834	13	4991
14	0557	14	3729	14	4838	14	4992
15	0596	15	3749	15	4842	15	4992
16	0636	16	3770	16	4846	16	4992
17	0675	17	3790	17	4850	17	4992
18	0714	18	3810	18	4854	18	4993
19	0753	19	3830	19	4857	19	4993
0,20	0,0793	1,20	0,3849	2,20	0,4861	3,20	0,4993
21	0832	21	3869	21	4864	21	4993
22	0871	22	3883	22	4868	22	4994
23	0910	23	3907	23	4871	23	4994
24	0948	24	3925	24	4875	24	4994
25	0987	25	3944	25	4878	25	4994
26	1026	26	3962	26	4881	26	4995
27	1064	27	3980	27	4884	27	4995
28	1103	28	3997	28	4887	28	4995
29	1141	29	4015	29	4890	29	4995
0,30	0,1179	1,30	0,4032	2,30	0,4893	3,30	0,4995
31	1217	31	4049	31	4896	31	4995
32	1255	32	4066	32	4898	32	4996
33	1293	33	4082	33	4901	33	4996
0,34	0,1331	1,34	0,4099	2,34	0,4904	3,34	0,4996
35	1368	35	4115	35	4906	35	4996
36	1406	36	4131	36	4909	36	4996
37	1443	37	4147	37	4911	37	4996

Продолжение таблицы E

x	$\Phi(x)$	x	$\Phi(x)$	x	$\Phi(x)$	x	$\Phi(x)$
38	1480	38	4162	38	4913	38	4996
39	1517	39	4177	39	4916	39	4997
0,40	0,1554	1,40	0,4192	2,40	0,4918	3,40	0,4997
41	1591	41	4207	41	4920	41	4997
42	1628	42	4222	42	4922	42	4997
43	1664	43	4236	43	4925	43	4997
44	1700	44	4251	44	4927	44	4997
45	1736	45	4265	45	4929	45	4997
46	1772	46	4279	46	4931	46	4997
47	1808	47	4229	47	4933	47	4997
48	1844	48	4306	48	4934	48	4998
49	1879	49	4319	49	4936	49	4998
0,50	0,1915	1,50	0,4332	2,50	0,4938	3,50	0,4998
51	1950	51	4345	51	4939	51	4998
52	1985	52	4357	52	4941	52	4998
53	2019	53	4370	53	4943	53	4998
54	2054	54	4382	54	4945	54	4998
55	2088	55	4394	55	4946	55	4998
56	2123	56	4406	56	4948	56	4998
57	2157	57	4418	57	4949	57	4998
58	2190	58	4429	58	4951	58	4998
59	2224	59	4441	59	4952	59	4998
0,60	0,2257	1,60	0,4452	2,60	0,4953	3,60	0,4998
61	2291	61	4463	61	4955	61	4998
62	2324	62	4474	62	4956	62	4998
63	2357	63	4484	63	4957	63	4998
64	2389	64	4495	64	4959	64	4999
65	2422	65	4505	65	4960	65	4999
66	2454	66	4515	66	4961	66	4999
67	2486	67	4525	67	4962	67	4999
68	2517	68	4535	68	4963	68	4999
69	2549	69	4545	69	4964	69	4999
0,70	0,2580	1,70	0,4554	2,70	0,4965	3,70	0,4999
71	2611	71	4564	71	4966	71	4999
72	2642	72	4573	72	4967	72	4999
73	2673	73	4582	73	4968	73	4999
74	2703	74	4591	74	4969	74	4999
75	2734	75	4599	75	4970	75	4999
76	2764	76	4608	76	4971	76	4999
77	2794	77	4616	77	4972	77	4999
78	2823	78	4625	78	4973	78	4999
79	2852	79	4633	79	4973	79	4999

Продолжение таблицы Е

x	$\Phi(x)$	x	$\Phi(x)$	x	$\Phi(x)$	x	$\Phi(x)$
0,80	0,2881	1,80	0,4641	2,80	0,4974	3,80	0,4999
81	2910	81	4669	81	4975	81	4999
82	2939	82	4656	82	4976	82	4999
83	2967	83	4664	83	4986	83	4999
84	2995	84	4671	84	4977	84	4999
85	3023	85	4678	85	4978	85	5000
86	3051	86	4686	86	4979	86	5000
87	3078	87	4693	87	4979	87	5000
88	3106	88	4699	88	4980	88	5000
89	3133	89	4706	89	4981	89	5000
0,90	0,3159	1,90	0,4713	2,90	0,4981	3,90	0,5000
91	3186	91	4719	91	4982	91	5000
92	3212	92	4726	92	4982	92	5000
93	3238	93	4732	93	4983	93	5000
94	3264	94	4738	94	4984	94	5000
95	3289	95	4744	95	4984	95	5000
96	3315	96	4750	96	4985	96	5000
97	3340	97	4756	97	4985	97	5000
98	3365	98	4761	98	4986	98	5000
99	3389	99	4767	99	4986	99	5000

Таблица Ж

Критические значения F-критерия Фишера при уровне
значимости $\alpha = 0,05$

K1/K2	1	2	3	4	5	6	8	12	24	∞
1	161,45	199,50	215,72	224,57	230,17	233,97	238,89	243,91	249,04	254,32
2	18,51	19,00	19,16	19,25	19,30	19,33	19,37	19,41	19,45	19,50
3	10,13	9,55	9,28	9,12	9,01	8,94	8,84	8,74	8,64	8,53
4	7,71	6,94	6,59	6,39	6,26	6,16	6,04	5,91	5,77	5,63
5	6,61	5,79	5,41	5,19	5,05	4,95	4,82	4,68	4,53	4,36
6	5,99	5,14	4,76	4,53	4,39	4,28	4,15	4,00	3,84	3,67
7	5,59	4,74	4,35	4,12	3,97	3,87	3,73	3,57	3,41	3,23
8	5,32	4,46	4,07	3,84	3,69	3,58	3,44	3,28	3,12	2,93
9	5,12	4,26	3,86	3,63	3,48	3,37	3,23	3,07	2,90	2,71
10	4,96	4,10	3,71	3,48	3,33	3,22	3,07	2,91	2,74	2,54
11	4,84	3,98	3,59	3,36	3,20	3,09	2,95	2,79	2,61	2,40
12	4,75	3,88	3,49	3,26	3,11	3,00	2,85	2,69	2,50	2,30
13	4,67	3,80	3,41	3,18	3,02	2,92	2,77	2,60	2,42	2,21

Продолжение таблицы Ж

K1/K2	1	2	3	4	5	6	8	12	24	∞
14	4,60	3,74	3,34	3,11	2,96	2,85	2,70	2,53	2,35	2,13
15	4,54	3,68	3,29	3,06	2,90	2,79	2,64	2,48	2,29	2,07
16	4,49	3,63	3,24	3,01	2,85	2,74	2,59	2,42	2,24	2,01
17	4,45	3,59	3,20	2,96	2,81	2,70	2,55	2,38	2,19	1,96
18	4,41	3,55	3,16	2,93	2,77	2,66	2,51	2,34	2,15	1,92
19	4,38	3,52	3,13	2,90	2,74	2,63	2,48	2,31	2,11	1,88
20	4,35	3,49	3,10	2,87	2,71	2,60	2,45	2,28	2,08	1,84
21	4,32	3,47	3,07	2,84	2,68	2,57	2,42	2,25	2,05	1,81
22	4,30	3,44	3,05	2,82	2,66	2,55	2,40	2,23	2,03	1,78
23	4,28	3,42	3,03	2,80	2,64	2,53	2,38	2,20	2,00	1,76
24	4,26	3,40	3,01	2,78	2,62	2,51	2,36	2,18	1,98	1,73
25	4,24	3,38	2,99	2,76	2,60	2,49	2,34	2,16	1,96	1,71
26	4,22	3,37	2,98	2,74	2,59	2,47	2,32	2,15	1,95	1,69
27	4,21	3,35	2,96	2,73	2,57	2,46	2,30	2,13	1,93	1,67
28	4,20	3,34	2,95	2,71	2,56	2,44	2,29	2,12	1,91	1,65
29	4,18	3,33	2,93	2,70	2,54	2,43	2,28	2,10	1,90	1,64
30	4,17	3,32	2,92	2,69	2,53	2,42	2,27	2,09	1,89	1,62
35	4,12	3,26	2,87	2,64	2,48	2,37	2,22	2,04	1,83	1,57
40	4,08	3,23	2,84	2,61	2,45	2,34	2,18	2,00	1,79	1,51
45	4,06	3,21	2,81	2,58	2,42	2,31	2,15	1,97	1,76	1,48
50	4,03	3,18	2,79	2,56	2,40	2,29	2,13	1,95	1,74	1,44
60	4,00	3,15	2,76	2,52	2,37	2,25	2,10	1,92	1,70	1,39
70	3,98	3,13	2,74	2,50	2,35	2,23	2,07	1,89	1,67	1,35
80	3,96	3,11	2,72	2,49	2,33	2,21	2,06	1,88	1,65	1,31
90	3,95	3,10	2,71	2,47	2,32	2,20	2,04	1,86	1,64	1,28
100	3,94	3,09	2,70	2,46	2,30	2,19	2,03	1,85	1,63	1,26
125	3,92	3,07	2,68	2,44	2,29	2,17	2,01	1,83	1,60	1,21
150	3,90	3,06	2,66	2,43	2,27	2,16	2,00	1,82	1,59	1,18
200	3,89	3,04	2,65	2,42	2,26	2,14	1,98	1,80	1,57	1,14
300	3,87	3,03	2,64	2,41	2,25	2,13	1,97	1,79	1,55	1,10
400	3,86	3,02	2,63	2,40	2,24	2,12	1,96	1,78	1,54	1,07
500	3,86	3,01	2,62	2,39	2,23	2,11	1,96	1,77	1,54	1,06
1000	3,85	3,00	2,61	2,38	2,22	2,10	1,95	1,76	1,53	1,03
∞	3,84	2,99	2,60	2,37	2,21	2,09	1,94	1,75	1,52	1,00

ПРЕДМЕТНЫЙ УКАЗАТЕЛЬ

А

Агрегирование 43
Асимметрии коэффициент 123
Асимметрическое распределение
81, 82, 84
Асимметрия 81, 82, 84, 123

Б

Бимодальность 73

В

Валидность 39
Вариабельность
(скедастичность) 145, 146
Вариации размах 84
Вариація 83
Варьируемость 83
Величина
– случайная 43
– средняя 65

Г

Гартовская диаграмма 61
Генеральная совокупность 15
Гетероскедастичность 146
Гипотеза 115, 117, 131
– альтернативная 117, 119, 130
– нулевая (нуля) 116, 117, 122
Гистограмма 53
Гомоскедастичность 146
Граф (-а) 214
– вершина 215
– дуга
– ориентированный 214, 215
– ребро 215
График (и) 53

– качественных данных 60
– многозначный 63

Группа социальная 197
– малая 197, 198, 206
Группировка (совокупности
данных) 45
Гуттман(а) 159

Д

Дескриптивный 7
Дециль 75, 76
Диаграмма 60
– временная 63
– круговая (гартовская) 61
– полос 60
Диапазон 86
– интерквартильный 86
– промежуточный 86, 87
Дисперсия 83, 84, 89

И

Измерение(-я) 9, 11
– в социологии 14
– единица 28
– метод(ы) 25
= присуждения баллов 26
= равных интервалов 25
= суммарных оценок 26
= шкалограммный анализ 26
= экспертных оценок 25
Индекс(ы) 100, 101, 210
– групповые 216
– персональные 210
– психологической взаимности
216
– социометрического статуса
210

- эмоциональной
экспансивности 216
- Интервал(ы) 45
 - доверительный 38
 - закрытые 45
 - неравные 45, 55
 - открытые 45
 - равные 45
- Интервалов равных метод 25

К

- Квантили 73
- Квартили 75
- Классификация 49
 - перекрестная 50
- Ковариация 140
- Ковариация 144
- Корреляционная (-ое) 135
 - зависимость 135
 - матрица 148
 - отношение 185
 - поле 141
 - связь 136
 - таблица 146, 147
- Корреляция 135
 - линейная 171, 176
 - множественная 171, 191
 - парная 171
 - рангов 162
 - частная 191, 192
- Коэффициент
 - ассоциации (Юла) 155
 - вариации 93, 94
 - = качественной 95
 - взаимной сопряженности 140
 - Гуттмана (λ) 159
 - детерминации 176
 - конкордации 169

- контингенции (Фи) 157
- корреляции 139
 - = Пирсона 179
 - = ранговой 163
- регрессии 191, 192
- сопряженности 151
 - = Пирсона 151
 - = Крамера 153, 154
 - = Чупрова 153
- Критерий (статистический) 76, 121
 - непараметрический 121
 - нормирующий 76
 - параметрический 121
 - Пирсона 123
 - согласия 121
 - Стьюдента (Т-критерий) 128
 - Фишера (Ф-критерий) 132
- Кумулята 53, 54, 57

Л

- Лидер(-ство) (в группе) 207, 208

М

- Максимум (как характеристика
распределения) 64, 80
- Математика 3, 197
- Математические методы 5, 6
- Математическое ожидание 65
- Матрица «2×2» 155
- Медиана 64, 73, 197
- Микро-, мезо-, макроуровень
социального взаимодействия 10
- Минимум (как характеристика
распределения) 64, 65, 80
- Мода 64, 71, 197
- Морено Я. Л. теория 198, 199

Н

Надежность (измерения) 30,31
Наклон линии регрессии 171
Накопленная частота 48
Накопленных частот график 57
Нормальная кривая 111, 112, 121
Нормирование элементарное 96

О

Обоснованность (в социологическом измерении) 32
Обратимость (коэффициента) 178
Операционализация 15
Отклонение (от среднего) 87, 172
– необъясняемые остаточные 140
– объясняемые остаточные 139
– среднее квадратическое 89, 91, 92
– среднее линейное 88, 92
Относительные показатели вариации 93
Ошибка 13
– второго рода 116
– первого рода 116
– систематическая 31
– случайная 32

П

Переменная(-ые) 15
– дискретные 17
– качественные 16
– количественные 17
– непрерывные 18
– упорядоченные 166

Пирсона коэффициент 151
– корреляции 179, 182
– сопряженности 151
Погрешность 32
Подклассификация 102, 103
Полигон распределения 56
Правильность (в социологическом измерении) 32, 33
Признак 15, 43
– группировочный 45
Прогноз 160
Программа исследования 202
Промех (в социологическом измерении) 31
Пропорции 97
Процентные отношения 97

Р

Разностей метод (вычисления моды) 72, 73
Ранжирование 29
Распределение 109
– двумерное 138
– закон 109
– идеальное 113
– кривое 109
– нормальное 109
– одномерное 110
– скошенное 122, 123
Распределения ряд 43
– вариационный 43
– динамический 44
– дискретный 44
– интервальный 44
– первичный 43
Рассеивание 143

- Регрессия(-ии) 174
– линия 174
– множественная 191, 192
– нелинейная 184
– частная 191, 192

С

- Связь 136
– криволинейная 171
– корреляционная 136
– прямолинейная 171
– слабая 171
– тесная 171
Семантический дифференциал 27
Сигма 91, 181
Сила связи (между признаками) 137
Скедастичность (вариабельность) 145, 146
Скошенность 122
Сопряженность 135
Социограмма 217
– групповая 219
– дифференциальная 221
– индивидуальная 221
– конвенциональная 220
– мишень 220
Социоматрица 208, 209
Социометрический (-ая, -ое, -ии)
– исследование 200
– критерии 202, 203
– процедура 201, 202
= параметрическая 201, 203
= непараметрическая 204
– таблица 209, 210
Социометрия 198
Спирмена коэффициент
(ранговой корреляции) 163

- Сплоченность групповая 211, 213,
Среднего арифметического свойства 70
Среднее арифметическое 66, 197
– взвешенное 67
– для интервальных данных 68
– комбинированное 68
– простое 67
Среднее вероятностное 71
Среднее(-ие) (величины) 65, 80, 179, 180
Статистическая карта 62
Статистический вывод 115
Статистический критерий 117, 118
– двусторонний (двусторонней области) 119
– односторонний (для односторонней области) 119
Статистического критерия мощность 120
Степень(-и) 98, 99
Структура группы 211, 214

Т

- Таблица статистическая 52
– групповая 52
– комбинационная 53
– перекрестная
– простая 52
Таблицы статистической подлежащее 52
Таблицы статистической сказуемое 52
Теснота связи 139, 140

У

Уровень значимости 116

Устойчивость (в измерении) 32, 36

Ф

Фактор 138

Х

Характеристики исследуемой совокупности 64

– положения 64

– рассеяния 82

Хи-квадрат 123, 150, 169

Ц

Ценностно-ориентационное единство (ЦОЕ) группы 213

Центиль 75, 76

Ч

Частость 44

Частота 43

– абсолютная 44

– кумулятивная (накоплен.) 48

– наблюдаемая (эмпирич.) 126

– относительная 44

– теоретическая 125, 126

Ш

Ширина интервала 46

Шкала 19, 20

– Богардуса 25

– высокого типа 20

– интервальная

– качественная 23

– количественная (числовая) 24

– метрическая 23

– низкого типа 20

– номинальная (наименований, классификационная) 20, 21

– отношений (абсолютная) 23

– порядковая (ранговая) 22

Э

Эмпирический индикатор 15

Эта-квадрат 185, 186, 187

Ю

Юла коэффициент 155

УСЛОВНЫЕ ОБОЗНАЧЕНИЯ

- Σ – сумма
 W – абсолютная устойчивость шкалы
 N – объем генеральной совокупности
 n – объем выборочной совокупности
 n_i – частота (абсолютная)
 n_i^H – кумулятивная частота (накопленная)
 n_{mo} – частота модального интервала
 n^- – частота интервала, предшествующего модальному интервалу
 n^+ – частота последующего после модального интервала
 n_H – частота, накопленная до медианного интервала
 n_{me} – частота медианного интервала
 $n(x_i)$ – число индивидов, у которых признак X принимает значение x_i
 $n(y_j)$ – число индивидов, у которых признак Y принимает значение y_j
 $\tilde{n}_{ij}$ – ожидаемая частота, то есть та, которая могла бы быть получена в случае полного отсутствия связи между признаками X и Y (теоретическая частота)
 n_{ij} – число индивидов, у которых признак X принимает значение x_i , а признак Y принимает значение y_j
 $X, Y \dots$ – признаки (переменные)
 v_i – частота относительная (частость, доля, %%)
 v_i^H – кумулятивная частость (доли, %%)
 h – ширина интервала
 k – число интервалов
 $M[X]$ или $\bar{x}$ – среднее арифметическое
 Mo – мода
 Me – медиана
 Q_i – i -й квартиль ($i = \overline{1,3}$)
 D_i – i -й дециль ($i = \overline{1,9}$)

C_i – i -й центиль ($i = \overline{1, 99}$)
 G_N – среднее геометрическое
 R – размах вариации
 E – энтропия
 ε – энтропийная мера дисперсии
 D или σ^2 – дисперсия
 $\bar{d}$ – среднее линейное отклонение
 σ – среднее квадратическое отклонение
 σ^2 – выборочная дисперсия (оценка σ^2)
 V_d – коэффициент вариации для среднего линейного отклонения
 V_δ – коэффициент вариации для среднего квадратического отклонения
 V_q – коэффициент качественной вариации
 x_i – i -е значение признака (переменной)
 x_{ci} – средняя точка (середина) интервала
 x_{n1} – левая (нижняя) граница интервала
 x_{n2} – правая (верхняя) граница интервала
 x_0 – нижняя граница модального, медианного интервала
 $\bar{x}_i$ и $\bar{y}_i$ – условные средние (например в подгруппе)
 p – доверительная вероятность
 χ^2 – критерий Хи-квадрат Пирсона
 $\chi^2_{набл}$ – эмпирическое (реально полученное) значение Хи-квадрат
 $\chi^2_{кр}$ – критическое (табличное) значение Хи-квадрат
 $f(x)$ – функция плотности распределения
 df – число степеней свободы
 α – уровень значимости
 r – число строк таблицы сопряженности
 c – число столбцов таблицы сопряженности
 z – критические значения нормального распределения
 β – вероятность совершения ошибки второго рода
 K – статистический критерий
 $K_{набл}$ – наблюдаемое значение статистического критерия

$K_{кр}$ – критическое значение критерия
 t – значение Т-распределения Стьюдента
 ϕ – значение распределения Фишера
 H_0 – нулевая гипотеза
 H_1 – альтернативная гипотеза
 As – коэффициент асимметрии
 C – коэффициент средней квадратической сопряженности (Пирсона)
 T – коэффициент Чупрова
 T_c – коэффициент Крамера
 Q – коэффициент ассоциации (Юла) для таблиц 2×2
 Φ – коэффициент контингенции (Фи) для таблиц 2×2
 r_s – коэффициент ранговой корреляции Спирмена
 τ – коэффициент ранговой корреляции Кендэла
 r – коэффициент корреляции Пирсона
 d – коэффициент Сомерса
 $r_{01.2}$ – коэффициент частной корреляции
 $R_{0.123}$ – коэффициент множественной корреляции
 W – коэффициент конкордации
 η – корреляционное отношение
 Г.С.И. – групповые социометрические индексы
 П.С.И. – персональные социометрические индексы
 d – социометрическое ограничение (лимит выборов)
 $P(A)$ – вероятность случайного события социометрического выбора
 C_i – социометрический статус i -члена группы
 $R_i (+, -)$ – полученные i -членом группы положительные или отрицательные выборы (при проведении социометрии)
 $R-$ – отрицательные выборы, полученные i -членом группы (при проведении социометрии)
 E_j – эмоциональная экспансивность j -члена группы (при проведении социометрии)
 $R_j (+, -)$ – сделанные j -членом группы положительные или отрицательные выборы (при проведении социометрии)
 $C_{гр}$ – индекс сплоченности группы

A_g – индекс экспансивности группы

G_g – индекс взаимности в группе

————→ позитивный односторонний выбор (на социограмме)

↔ позитивный обоюдный выбор (на социограмме)

-----> негативный односторонний выбор (на социограмме)

↔-----> негативный обоюдный выбор (на социограмме)

СПИСОК ИСПОЛЬЗОВАННЫХ И РЕКОМЕНДУЕМЫХ ИСТОЧНИКОВ

Источники к ВВЕДЕНИЮ и РАЗДЕЛУ I

Основные источники

1. Грес П. В. Математика для гуманитариев : учеб. пособие / П. В. Грес. – М. : Юрайт, 2000. – 112 с.
2. Осипов Г. В. Методы измерения в социологии / Г. В. Осипов. – М. : Наука, 2003. – 124 с.
3. Паниотто В. И. Количественные методы в социологических исследованиях / В. И. Паниотто, В. С. Максименко. – К. : Наук. думка, 1982. – 272 с.
4. Паніотто В. І. Статистичний аналіз соціологічних даних / В. І. Паніотто, В. М. Максименко, Н. М. Харченко. – К. : КМ Академія, 2004. – 270 с.
5. Татарова Г. Г. Методология анализа данных в социологии : учебник для вузов [2-е издание] / Г. Г. Татарова. – М. : Note Bene, 1999. – 180 с.
6. Циба В. Т. Математичні основи соціологічних досліджень. Кваліметричний підхід / В. Т. Циба – К. : Персонал (МАУП), 2002. – 248 с.
7. Ядов В. А. Стратегия социологического исследования. Описание, объяснение, понимание социальной реальности / В. А. Ядов. – М. : Добросвет, 1999. – 596 с.
8. Яковенко Т. В. Прикладна статистика в соціології. Практикум : навч. посіб. / Т. В. Яковенко. – Х. : Вид-во ХНУ, 2003. – 238 с.

Дополнительные источники

9. Амберкромби Н. Социологический словарь : [пер. с англ.] / Н. Аберкромби, С. Хилл, Б. С. Тернер. – М. : ЗАО «Издательство «Экономика», 2004. – 509 с.
10. Андреенков В. Г. Анализ и интерпретация эмпирических данных // Социология: Основы общей теории : учеб. пособ. для

высш. учеб. заведений / [Г. В. Осипов, Л. Н. Москвичев, А. В. Кабыща и др. : редколлегия Г. В. Осипов (отв. ред.), Л. Н. Москвичев (отв. ред.) и др.]. – М. : Аспект Пресс, 1996. – 460 с.

11. Берка К. Измерения: понятия, теории, проблемы / К. Берка. – М. : Прогресс, 1987. – 210 с.

12. Большой энциклопедический словарь [Электронный ресурс]. – Режим доступа: <http://www.slovopedia.com> .

13. Википедия (свободная энциклопедия). – Режим доступа: <http://www.wikipedia.org> .

14. Зимняя И. А. Ключевые компетенции – новая парадигма результатов образования / И. А. Зимняя // Высшее образование сегодня. – 2003. – № 5. – С. 34–42.

15. Измерение в социологии : социологический словарь [Электронный ресурс]. – Режим доступа: http://mirslovarei.com/content_soc.

16. Орлов А. И. Прикладная статистика / А. И. Орлов. – М. : Экзамен, 2004. – 320 с.

17. Осипов Г. В. Рабочая книга социолога [Изд. 4] / Г. В. Осипов. – М. : КомКнига, 2006. – 480 с.

18. Паниотто В. И. Качество социологической информации. – К. : Наук. думка, 1986. – 206 с.

19. Попов О. А. Измерение в психологии. Шкалы измерения [Электронный ресурс] / О. А. Попов. – Режим доступа: <http://psystat.at.ua/publ/1-1-0-28> .

20. Ротман Д. Г. Оперативные социологические исследования / Д. Г. Ротман, С. Н. Бурова, Н. П. Веремеева и др. – Мн. : Веды, 1997. – 208 с.

21. Суппес П. Основы теории измерений // Психологические измерения / П. Суппес, Дж. Зинес. – М. : Мир, 1967. – 287 с.

22. Пфанцагль И. Теория измерений / И. Пфанцагль. – М., 1976. – 250 с.

23. Саганенко Г. И. Уровни эмпирической доказательности в социологии / Г. И. Саганенко. – М. : ИС АН СССР, 1991. – 74 с.

24. Саганенко Г. И. Надежность результатов социологического исследования / Г. И. Саганенко. – Л. : Наука, 1983.

25. Толстова Ю. Н. Измерение в социологии : курс лекций / Ю. Н. Толстова. – М. : Инфра-М, 1998. – 224 с.
26. Толстова Ю. Н. Кризис социологического измерения в начале нашего века и пути выхода из него / Ю. Н. Толстова // Социология. – 1995. – № 5/6. – С. 103–117.
27. Хуторской А. В. Ключевые компетенции и образовательные стандарты: Доклад на отделении философии образования и теории педагогики РАО, 23 апреля 2002 (Центр «Эйдос») [Электронный ресурс] / А. В. Хуторской. – Режим доступа: [www/eidos.ru/news/compnet/htm](http://eidos.ru/news/compnet/htm).
28. Шереги Ф. Э. Основы прикладной социологии / Ф. Э. Шереги, М. К. Горшков. – М. : Интерпракс, 1996. – 184 с.
29. Циба В. Т. Основи теорії кваліметрії : навч. посіб. / В. Т. Циба. – К. : ІЗМН, 1997. – 160 с.
30. Яковенко Т. В. Прикладна статистика в соціології. Практичні завдання. Навчально-методична розробка для студентів соціологічного факультету / Т. В. Яковенко. – Х. : Вид-во Харк. держ. ун-ту харчування та торгівлі, 2006. – 92 с.
31. Ястребов А. В. Междисциплинарный подход в преподавании математики / А. В. Ястребов // Ярославский педагогический вестник. – 2004. – № 3 (40). – С. 5–15.
32. Campbell N. R. Foundations of science/The Philosophy of Theory and Experiment / N. R. Campbell. – N.Y., Dover, 1957.
33. Krantz D. H. Foundations of measurement / D. H. Krantz, R. D. Luce, P. Suppes, A. Tversky. – N.Y., L., 1971.

Источники к РАЗДЕЛУ II

Основные источники

1. Борисова Е. В. Формирование и математическая обработка данных в социологии : учеб. пособие ; [1-е изд.] / Е. В. Борисова. – Тверь : ТГТУ, 2006. – 120 с.
2. Васнев С. А. Статистика : учеб. пособие / С. А. Васнев. – М. : МГУП, 2001. – 170 с.

3. Гмурман В. Е. Руководство к решению задач по теории вероятностей и математической статистике / В. Е. Гмурман. – М. : Высш. шк., 2004. – 404 с.
4. Паниотто В. И. Количественные методы в социологических исследованиях / В. И. Паниотто, В. С. Максименко. – К. : Наук. думка, 1982. – 272 с.
5. Паніотто В. І. Статистичний аналіз соціологічних даних / В. І. Паніотто, В. М. Максименко, Н. М. Харченко. – К. : КМ Академія, 2004. – 270 с.
6. Осипов Г. В. Рабочая книга социолога [Изд.4] / Г. В. Осипов. – М. : КомКнига, 2006. – 480 с.
7. Телейко А. Б. Математико-статистичні методи в соціології та психології / А. Б. Телейко, Р. К. Чорней. – К. : МАУП, 2007. – 424 с.
8. Циба В. Т. Математичні основи соціологічних досліджень. Кваліметричний підхід / В. Т. Циба – К. : Персонал (МАУП), 2002. – 248 с.
9. Яковенко Т. В. Прикладна статистика в соціології. Практикум : навч. посіб. / Т. В. Яковенко. – Х. : Вид-во ХНУ, 2003. – 238 с.

Дополнительные источники

10. Вентцель Е. С. Теория вероятностей / Е. С. Вентцель. – М. : Наука, 1969. – 576 с.
11. Ганнушкина С. А. Сборник задач по теории вероятностей / С. А. Ганнушкина, В. Ю. Сеницын. – М. : РГГУ, 1997. – 52 с.
12. Гмурман В. Е. Теория вероятностей и математическая статистика / В. Е. Гмурман. – М. : Высш. шк., 2003. – 480 с.
13. Елисеева И. И. Общая теория статистики / И. И. Елисеева, М. М. Юзбашев. – М. : Финансы и статистика, 1995. – 365 с.
14. Калинина В. Н. Математическая статистика / В. Н. Калинина, В. Ф. Панкин. – М. : Дрофа, 2002. – 336 с.
15. Кремер Н. Ш. Теория вероятностей и математическая статистика / Н. Ш. Кремер. – М. : ЮНИТИ, 2007. – 551 с.
16. Курс социально-экономической статистики : учебник для

вузов ; [под ред. М. Г. Назарова]. – М. : Финстатинформ, Юнити-Дана, 2000. – 230 с.

17. Письменный Д. Т. Конспект лекций по теории вероятностей / Д. Т. Письменный. – М. : Айрис-пресс, 2008. – 256 с.

18. Шведов А. С. Теория вероятностей и математическая статистика [уч. пособ. для экон. спец. вузов]. – М. : Изд-во Высшей школы экономики, 1995. – 208 с.

19. Шмойлова Р. А. Теория статистики / Р. А. Шмойлова. – М. : Финансы и статистика, 2000. – 558 с.

20. Эддоус М. Методы принятия решений / М. Эддоус, Р. Стэнсфилд. – М. : Аудит; ЮНИТИ, 1997. – 590 с.

21. Электронный учебник по статистике [электронный ресурс] : Информ.-аналит. материалы. – М., StatSoft. – Режим доступа: <http://www.statsoft.ru/home/textbook/default.htm>

Источники к РАЗДЕЛУ III

Основные источники

1. Елисеева И. И. Общая теория статистики : учебник / И. И. Елисеева, М. М. Юзбашев ; [под ред. чл.-корр. РАН И. И. Елисеевой]. – 4-е изд., перераб. и доп., – М. : Финансы и статистика, 2001. – 480 с.

2. Орлов А. И. Математика случая: вероятность и статистика – основные факты : учебное пособие / А. И. Орлов. – М. : МЗ-Пресс, 2004. – 250 с.

3. Сечко В. В. Математические методы обработки психологических данных / В. В. Сечко. – Минск, 2002. – 80 с.

4. Паниотто В. И. Количественные методы в социологических исследованиях / В. И. Паниотто, В. С. Максименко. – К. : Наук. думка, 1982. – 272 с.

5. Паніотто В. І. Статистичний аналіз соціологічних даних / В. І. Паніотто, В. М. Максименко, Н. М. Харченко. – К. : КМ Академія, 2004. – 270 с.

6. Ядов В. А. Стратегия социологического исследования.

Описание, объяснение, понимание социальной реальности / В. А. Ядов. – М. : Добросвет, 1999. – 596 с.

7. Телейко А. Б. Математико-статистичні методи в соціології та психології / А. Б. Телейко, Р. К. Чорней. – К. : МАУП, 2007. – 424 с.

8. Толстова Ю. Н. Анализ социологических данных. Методология, дескриптивная статистика, изучение связей между номинальными признаками / Ю. Н. Толстова. – М. : Научный мир, 2000. – 352 с.

9. Циба В. Т. Математичні основи соціологічних досліджень. Кваліметричний підхід / В. Т. Циба – К. : Персонал (МАУП), 2002. – 248 с.

10. Яковенко Т. В. Прикладна статистика в соціології. Практикум : навч. посіб. / Т. В. Яковенко. – Х. : Вид-во ХНУ, 2003. – 238 с.

Дополнительные источники

11. Андриашин Х. А. Информатика и математика для юристов (Рекомендовано МО РФ в качестве учебного пособия для студентов ВУЗов, обучающихся по юридическим специальностям) / Х. А. Андриашин, С. Я. Казанцев, О. Э. Згадзай. – М. : ЮНИТИ, 2003. – 463 с.

12. Барвенков С. А. Лекции по высшей математике [Электронной издание] / С. А. Барвенков. – Режим доступа: <http://www.kursach.com/urista.htm>.

13. Кобзарь А. И. Прикладная математическая статистика. Для инженеров и научных работников / А. И. Кобзарь – М. : Физматлит, 2006. – 816 с.

14. Лемешко Б. Ю. Сравнительный анализ критериев проверки отклонения распределения от нормального закона / Б. Ю. Лемешко, С. Б. Лемешко // Метрология. – 2005. – № 2. – С.3–23.

15. Рассолов М. М. Элементы высшей математики для юристов : учеб. пособие / М. М. Рассолов, С. Г. Чубукова, В. Д. Элькин. – М. : ЮРИСТЪ, 1999. – 184 с.

16. Титкова Л. С. Математические методы в психологии /

Л. С. Титкова. – Владивосток : изд-во Дальневосточного университета, 2002. – 150 с.

17. Центр Разумкова. Официальный сайт [Электронный ресурс]. – Режим доступа: http://www.uceps.org/ukr/expert.php?news_id=3293

18. Шеломовский В. В. Математическая статистика [Электронный ресурс] / В. В. Шеломовский. – МФГПУ. – Режим доступа: <http://www.exponenta.ru/educat/systemat/shelomovsky/lab/lab15.asp>

Источники к РАЗДЕЛУ IV

Основные источники

1. Бессокирная Г. П. Дискриминантный анализ / Г. П. Бессокирная // Социология: 4М. – 2003. – №16. – С. 25-35.

2. Гаптон А. Анализ таблиц сопряженности. [Электронная библиотека МГУ им. М. В. Ломоносова]. – Режим доступа: <http://lib.socio.msu.ru/l/library>

3. Девятко И. Ф. Методы социологического исследования / И. Ф. Девятко. – Екатеринбург : изд-во Уральск. ун-та, 1998 г. – 208 с.

4. Наследов А. Д. Математические методы психологического исследования. Анализ и интерпретация данных : учеб. пособие; 4-е изд. / А. Д. Наследов. – СПб. : Речь, 2012. – 392 с.

5. Паниотто В. И. Количественные методы в социологических исследованиях / В. И. Паниотто, В. С. Максименко. – К. : Наук. думка, 1982. – 272 с.

6. Паніотто В. І. Статистичний аналіз соціологічних даних / В. І. Паніотто, В. М. Максименко, Н. М. Харченко. – К. : КМ Академія, 2004. – 270 с.

7. Телейко А. Б. Математико-статистичні методи в соціології та психології / А. Б. Телейко, Р. К. Чорней. – К. : МАУП, 2007. – 424 с.

8. Толстова Ю. Н. Анализ социологических данных. Методология, дескриптивная статистика, изучение связей между

номинальными признаками / Ю. Н. Толстова. – М. : Научный мир, 2000. – 352 с.

9. Яковенко Т. В. Прикладна статистика в соціології. Практикум : навч. посіб. / Т. В. Яковенко. – Х. : Вид-во ХНУ, 2003. – 238 с.

Дополнительные источники

10. Ибрагимов Н. М. Регрессионный анализ / Н. М. Ибрагимов, В. В. Карпенко, Е. А. Коломак, В. И. Суслов. – ТЕМПУС, Экономический факультет НГУ, 1997. – 120 с.

11. Ключенко Э. Политическое участие: теория, методология и измерение с применением метода шкалограммирования по Гуттману / Э. Ключенко // Социология: теория, методы, маркетинг. – 2005. – №4. – С. 46–72.

12. Колесов В. М. Коэффициенты связи для совокупности номинальных признаков / В. М. Колесов // Социология: 4М. – Т. 1. – №1. – С. 62–74.

13. Пак Т. В. Эконометрика : учеб. пособие / Т. В. Пак, Я. И. Еремеева. – Владивосток : Изд-во Дальневост. ун-та, 2009. – 70 с.

14. Титкова Л. С. Математические методы в психологии / Л. С. Титкова. – Владивосток : изд-во Дальневосточного университета, 2002. – 150 с.

15. Толстова Ю. Н. Измерение в социологии : курс лекций / Ю. Н. Толстова. – М. : ИНФРА-М, 1998. – 224 с.

16. Циба В. Т. Математичні основи соціологічних досліджень. Кваліметричний підхід / В. Т. Циба – К. : Персонал (МАУП), 2002. – 248 с.

17. Чесноков С. В. Основы гуманитарных измерений / С. В. Чесноков. – М. : ВНИИСИ, 1986. – С. 20–90.

18. Яковенко Т. В. Прикладна статистика в соціології. Практичні завдання. Навчально-методична розробка для студентів соціологічного факультету / Т. В. Яковенко. – Х. : Вид-во Харк. держ. ун-ту харчування та торгівлі, 2006. – 92 с.

Источники к РАЗДЕЛУ V

Основные источники

1. Бурков В. Н. Теория графов в управлении организационными системами / В. Н. Бурков, А. Ю. Заложнев, Д. А. Новиков. – М. : Синтег, 2001. – 124 с.
2. Душков Б. А. Энциклопедический словарь: психология труда, управления, инженерная психология и эргономика [Электронный ресурс] / Б. А. Душков, А. В. Королев, Б. А. Смирнов. – 2005. – Режим доступа: <http://vocabulary.ru/dictionary/896/word>.
3. Осипов Г. В. Методы измерения в социологии / Г. В. Осипов. – М. : Наука, 2003. – 124 с.
4. Паниотто В. И. Структура межличностных отношений. Методика и математические методы исследования / В. И. Паниотто. – К. : Наук. думка, 1975. – 128 с.
5. Паниотто В. И. Количественные методы в социологических исследованиях / В. И. Паниотто, В. С. Максименко. – К. : Наук. думка, 1982. – 272 с.
6. Паниотто В. И. Статистичний аналіз соціологічних даних / В. І. Паніотто, В. М. Максименко, Н. М. Харченко. – К. : КМ Академія, 2004. – 270 с.
7. Социометрия: исследование межличностных отношений в группе [Электронный ресурс]. – Режим доступа: <http://psyfactor.org/moreno.htm>.

Дополнительные источники

8. Андреева Г. М. Социальная психология [Электронный ресурс] / Г. М. Андреева. – М. : Наука, 1994. – Режим доступа: <http://psylib.org.ua/books/andrg01/index.htm>.
9. Бурков В. Н. Прикладные задачи теории графов / В. Н. Бурков, И. А. Горгидзе, С. Е. Ловецкий. – Тбилиси : Мецниереба, 1974. – 234 с.
10. Десев Л. Психология малых групп: Социальные иллюзии и проблемы / Л. Досев ; [пер. с болг. Т. Е. Зюзкжиной. Общ. ред. и послесл. Н. Ф. Наумовой]. – М. : Прогресс, 1979. – 210 с.

11. Донцов А. И. Проблема групповой сплоченности / А. И. Донцов. – М. : МГУ, 1979. – 128 с.
12. Донцов А. И. О понятии группа в социальной психологии / И. А. Донцов // Вестник Моск. ун.-та. Сер. 14. Психология. – 1997. – № 4. – С.17–25.
13. Донцов А. И. Психология коллектива / А. И. Донцов. – М. : МГУ, 1984. –106 с.
14. Золотовицкий Р. Социометрия: мера общения / Р. Золотовицкий // Флогистон [Электронный ресурс]. – 2002, 20 апреля. – Режим доступа: http://flogiston.ru/articles/social/moreno_zolotovitsky.
15. Метод социометрических измерений // Энциклопедия психодиагностики [Электронный ресурс]. – Режим доступа: <http://psylab.info>.
16. Морено Я. Социометрия [Электронный ресурс] / Я. Морено. – Режим доступа: http://polbu.ru/moreno_sociometry.
17. Немов Р. С. Психология : учеб. пособие для учащихся пед. уч-щ, студентов пед. ин-тов и работников системы подготовки, повышения квалификации и переподготовки пед. кадров / Р. С. Немов. – М. : Просвещение, 1990. – 301 с.
18. Паниотто В. И. Социометрические методы изучения малых социальных групп / В. И. Паниотто // Социол. исслед. – 1976. – №3. – С. 152–156.

СОДЕРЖАНИЕ

Введение	3
Раздел I. Измерение как предпосылка применения математических методов в социологии	9
1.1. Измерения в социологии: теоретико-методологический ракурс	9
1.2. Совокупности, признаки, переменные: к вопросу о том: «что измеряется в социологии?»	19
1.3. Шкалы, типы шкал, возможные операции со шкалами: к вопросу о том: «как происходит измерение в социологии?»	24
1.4. Надежность измерения	39
<i>Вопросы и задания для самоконтроля</i>	<i>53</i>
Раздел II. Агрегирование социологических данных и соответствующие математические процедуры	55
2.1. Частотные распределения	55
2.2. Перекрестная классификация и табличное представление социологической информации	64
2.3. Графическое представление социологической информации	69
2.4. Характеристики положения (среднее арифметическое, мода, медиана, квантили)	81
2.5. Критерии выбора вида усреднения	100
2.6. Характеристики рассеяния (дисперсия, отклонение, коэффициент вариации)	105
2.7. Вариация качественных переменных	120
<i>Вопросы и задания для самоконтроля</i>	<i>134</i>
Раздел III. Понятия и принципы статистического вывода как основа анализа социологических данных	140
3.1. Кривые распределения, законы распределения	140
3.2. Проверка нормальности распределения	147
3.3. Общие представления о критериях «нормальности» распределения	151

3.4. Критерии «нормальности»:	
принципы нахождения и расчетные формулы	156
<i>Вопросы и задания для самоконтроля</i>	<i>171</i>

Раздел IV. Измерение связи между признаками с помощью математических методов 173

4.1. Понятие статистической связи	173
4.2. Корреляционное поле	182
4.3. Корреляционная таблица	188
4.4. Коэффициенты связи, основанные на Хи-квадрат Пирсона (для номинальных признаков)	192
4.5. Коэффициенты связи для матриц 2г2 (для номинальных и порядковых признаков)	199
4.6. Коэффициенты связи, основанные на моделях прогноза (для номинальных признаков)	204
4.7. Корреляция рангов	208
4.8. Линейная корреляция (для метрических признаков и интервального уровня измерения)	218
4.9. Нелинейная регрессия. Множественная и частная корреляция	236
<i>Вопросы и задания для самоконтроля</i>	<i>250</i>

Раздел V. Применение математических методов при исследовании малых социальных групп 252

5.1. Социометрический опрос как метод социологического исследования малых социальных групп: специфика и общая схема действий	252
5.2. Общая процедура проведения социометрического исследования	259
5.3. Изучение внутригрупповой структуры, выявление «социометрических позиций» членов группы. Построение социоматрицы и вычисление персональных социометрических индексов	264

5.4. Изучение степени сплоченности/разобщенности группы. Примеры вычисления групповых социометрических индексов	270
5.5. Обнаружение внутригрупповых подсистем по резуль- татам социометрического исследования. Построение социограмм с обращением к теории графов	278
<i>Вопросы для самоконтроля</i>	285
Приложения	287
Предметный указатель	298
Условные обозначения	303
Список использованных и рекомендуемых источников ...	307

Навчальне видання

НЕЧИТАЙЛО Ірина Сергіївна
БІРЮКОВА Марина Василівна

МАТЕМАТИЧНІ МЕТОДИ В СОЦІОЛОГІЇ

Підручник для студентів вищих навчальних закладів

(російською мовою)

Редактор *З. М. Москаленко*
Комп'ютерна верстка *І. С. Кордюк*
Комп'ютерний набір *І. С. Нечитайло*
Художник *Р. М. Третьяков*

Підписано до друку 27.11.2013. Формат 60×84/16.

Папір офсетний. Гарнітура «Таймс».

Ум. друк. арк. 18,60. Обл.-вид. арк. 16,60.

Тираж 16 пр. Зам. №

План 2011/12 навч. р., поз. № 14 в переліку робіт кафедри

Видавництво
Народної української академії
Свідоцтво № 1153 від 16.12.2002.

Надруковано у видавництві
Народної української академії

Україна, 61000, Харків, МСП, вул. Лермонтовська, 27.