

С.М. ЕРМАКОВ
А.А. ЖИГЛЯВСКИЙ

МАТЕМАТИЧЕСКАЯ ТЕОРИЯ ОПТИМАЛЬНОГО ЭКСПЕРИМЕНТА



Е-42.
С. М. ЕРМАКОВ
А. А. ЖИГЛЯВСКИЙ

МАТЕМАТИЧЕСКАЯ ТЕОРИЯ ОПТИМАЛЬНОГО ЭКСПЕРИМЕНТА

С предисловием Г. И. МАРЧУКА

Допущено Министерством

высшей и средней специальной образования СССР

как качественное учебное пособие для студентов вузов

специальности «Прикладная математика»



МОСКВА «НАУКА»
ГЛАВНАЯ РЕДАКЦИЯ
ФИЗИКО-МАТЕМАТИЧЕСКОЙ ЛИТЕРАТУРЫ
1987



ББК 22.18

Е 72

УДК 519.6

Ермаков С. М., Жиглявский А. А. Математическая теория оптимального эксперимента: Учеб. пособие.— М. Наука. Гл. ред. физ.-мат. лит., 1987. — 320 с.

В основу положен курс лекций по оптимальному планированию эксперимента в ЛГУ. Систематически излагается математическая теория планирования эксперимента. Имеется много примеров и упражнений. Включен не только классический материал, но и новейшие результаты.

Для студентов факультетов и отделений прикладной математики, аспирантов и преподавателей вузов, а также для инженеров.

Табл. 9. Ил. 7. Библиогр. 52 назв.

Рецензенты:

кафедра автоматики МЭИ

доктор физико-математических наук М. Б. Малютов

ОГЛАВЛЕНИЕ

Предисловие	5
Предисловие авторов	6
Введение	9
Глава 1. Линейный регрессионный анализ	13
§ 1. Классическая линейная регрессионная модель и ее обобщения	13
§ 2. Линейный регрессионный анализ при наличии априорной информации о параметрах	41
§ 3. Основные понятия дисперсионного анализа	45
Глава 2. Факторные планы	51
§ 1. Полные факторные планы и их дробные реплики	53
§ 2. Латинские планы	62
§ 3. Неполноблочные планы	70
Глава 3. Классическая теория планирования эксперимента по оцениванию параметров линейных регрессионных моделей	83
§ 1. Основные понятия и вспомогательные результаты	83
§ 2. Теоремы эквивалентности	105
§ 3. Некоторые следствия теорем эквивалентности	114
§ 4. Численные методы построения оптимальных планов	126
Глава 4. Планирование регрессионного эксперимента (дальнейшие постановки задач и результаты)	139
§ 1. Планы первого порядка	139
§ 2. Некоторые обобщения классической постановки задачи планирования регрессионного эксперимента	148
§ 3. Линейная теория возмущений и планирование эксперимента	166
§ 4. Планирование эксперимента при неадекватности линейной модели	171
Глава 5. Анализ и планирование эксперимента для нелинейных регрессионных моделей	192
§ 1. Нелинейный регрессионный анализ	192
§ 2. Планирование эксперимента по оцениванию параметров нелинейной регрессии	200
§ 3. Последовательное планирование	204

Глава 6. Планирование экстремального эксперимента	210
§ 1. Сходимость псевдоградиентных алгоритмов	210
§ 2. Методы планирования экстремального эксперимента при наличии ограничений	218
§ 3. Поисковые алгоритмы	223
§ 4. Выбор длины шага и направления движения в методе крутого восхождения	235
Глава 7. Планирование эксперимента по проверке гипотез	247
§ 1. Планирование дискриминирующих экспериментов	247
§ 2. Планирование отсеивающих экспериментов	253
Глава 8. Планирование имитационного эксперимента	271
§ 1. Имитационный эксперимент	271
§ 2. Метод существенной выборки при одновременном оценивании нескольких интегралов	284
Приложение 1. Сведения из теории матриц	298
§ 1. Матричная алгебра	298
§ 2. Неравенства выпуклости	305
§ 3. Матричный анализ	308
Приложение 2. Нормальное распределение и распределения, связан- ные с ним	312
Список литературы	315
Предметный указатель	317

ПРЕДИСЛОВИЕ

В связи с многочисленными приложениями научного и технического характера получила широкое развитие теория планирования эксперимента, которая развивается на стыке таких математических дисциплин как вычислительная математика, математическая статистика, дискретная математика, теория оптимизации. Число работ (как теоретических, так и прикладных) в области планирования эксперимента практически необозримо. В этой связи актуальной задачей является создание соответствующих учебных руководств.

Предлагаемый вниманию читателей учебник в доступной форме и на хорошем математическом уровне позволяет ознакомиться с основными направлениями современной теории планирования эксперимента. Наряду с традиционными разделами он содержит такие важные для развития науки разделы, как планирование эксперимента для решения обратных задач (основанное на теории возмущений), связанный с ним раздел, посвященный планированию эксперимента в функциональных пространствах, а также теорию имитационного и отсеивающего экспериментов. Впервые в учебной литературе освещаются вопросы, в которые авторы внесли значительный личный вклад: планирование эксперимента при неадекватности линейной модели, планирование имитационного и экстремального экспериментов.

Сказанное позволяет рекомендовать книгу С. М. Ермакова, А. А. Жиглявского «Математическая теория оптимального эксперимента» широкому кругу читателей — студентам, инженерам и научным работникам.

Академик *Г. И. Марчук*

ПРЕДИСЛОВИЕ АВТОРОВ

Книга написана на основе лекционных курсов, читавшихся авторами в течение ряда лет на математико-механическом факультете Ленинградского государственного университета им. А. А. Жданова.

Методы планирования эксперимента, возникшие первоначально в связи со статистическими приложениями в сельском хозяйстве и медицине, получили в последние десятилетия широкое развитие и новые области применения. Особое значение эти методы приобрели в связи с крупными программами проведения и автоматизации научных исследований. Известно, что применение методов планирования эксперимента может дать более или менее значительный экономический эффект, но отсутствие соответствующего плана может сделать экспериментальную программу полностью безрезультатной. Это диктует необходимость подготовки специалистов различного уровня, владеющих уже известными методами планирования эксперимента и способных разрабатывать новые методы применительно к различным предметным областям.

Книга ориентирована на студентов факультетов и отделений прикладной математики, математиков и инженеров с повышенной математической подготовкой. Наряду с подробным изложением разделов, ставших ныне традиционными (теория D -оптимальных непрерывных планов регрессионного эксперимента, элементы теории блок-схем и др.), курс включает новые разделы, вошедшие в теорию оптимального эксперимента в последнее десятилетие (робастное, или несмещенное, планирование, планирование эксперимента в функциональных пространствах, планирование имитационного эксперимента). В книге сделана попытка проследить также связи между различными разделами теории.

Важный вопрос, который остался вне рамок книги — это создание математической модели реального эксперимента. Переход к такой модели должен осуществляться в рамках соответствующей предметной области.

Для сельского хозяйства подобный переход требует понимания аграрных или животноводческих вопросов, в химии нужно иметь априорное понимание условий протекания соответствующей

щих реакций: аналогичные проблемы возникают в физике, технике, биологии и др. В каждой области существуют специфические ограничения, и только специалисты могут указать группы факторов, которыми можно пренебречь. Вместе с тем специалист в конкретной предметной области редко владеет в полной мере математическим аппаратом. Обычно планирование реального эксперимента требует совместной работы группы специалистов. Будущим математикам — участникам такой группы — в первую очередь адресована эта книга.

Хотя планирование эксперимента, в особенности с использованием ЭВМ, может дать большой экономический эффект, далеко не всегда бывает просто преодолеть дополнительные трудности, связанные с организацией планирования эксперимента. Оно требует достаточно точного осуществления условий оптимального протекания эксперимента, что в свою очередь может привести к дополнительным затратам и предъявляет повышенные требования к квалификации персонала, осуществляющего эксперимент. Эти соображения должны, безусловно, иметь в виду каждый специалист в области математической теории планирования эксперимента.

Теория планирования эксперимента — наука сравнительно молодая, сформировавшаяся немногим более полувека назад. Основателем теории принято считать английского статистика Р. Фишера, а основополагающей работой — его книгу [49]. Ряд работ Фишера посвящен разработке методов статистического анализа и планирования эксперимента в задачах повышения урожайности сельскохозяйственной продукции. Из этих работ впоследствии развились дисперсионный анализ и факторное планирование.

Наличие большого числа прикладных задач стимулировало развитие математической теории оптимального эксперимента. Среди зарубежных ученых, внесших большой вклад в ее развитие, прежде всего нужно назвать американских ученых Дж. Бокса и Дж. Кифера, которые в 50-х годах опубликовали несколько выдающихся результатов. Боксу принадлежат результаты по планированию регрессионного эксперимента для полиномиальных функций регрессии, планированию экстремальных экспериментов, конструированию критериев оптимальности. Кифер в значительной мере является создателем классической теории планирования регрессионного эксперимента — наиболее развитого в настоящее время раздела теории оптимального эксперимента.

Ряд оригинальных направлений в теории планирования эксперимента возник в нашей стране благодаря влиянию советской школы теории вероятностей и математической статистики (А. Н. Колмогоров, Ю. В. Линник и их ученики) с одной стороны, и вычислительной математики (Г. И. Марчук и его ученики) — с другой. Существенный вклад в развитие теории планирования эксперимента и ее приложений внесли советские

ученые В. В. Налимов, В. В. Федоров, Г. К. Круг, Е. В. Маркова, М. Б. Малютов и другие. Среди работ, оказавших большое влияние на развитие теории имитационного эксперимента, можно отметить работы Н. Н. Ченцова и Г. А. Михайлова.

С современным состоянием теории планирования эксперимента можно ознакомиться по [16, 46, 48, 51]. Дополнительные сведения по регрессионному анализу можно получить в [11, 30, 35, 36, 52], по дисперсионному анализу — в [35, 44], по теории матриц — в [5, 9, 35]; по планированию эксперимента: факторного — в [8, 26, 39, 43], регрессионного — в [4, 12, 15, 23, 24, 25, 40, 41, 47, 48, 50], экстремального — в [1, 10, 18, 32, 33], дискриминирующего — в [40], имитационного — в [14, 17, 18, 19, 29]. Работы [1, 3, 10, 13, 19, 22, 23, 26, 30, 31, 38, 47] доступны лицам, не обладающим высокой математической подготовкой.

Для понимания материала книги от читателя требуется владение аппаратом теории матриц в объеме большем, чем это содержится в стандартных учебниках по алгебре (дополнительные сведения из теории матриц изложены в приложении 1) и основными понятиями теории вероятностей и математической статистики (в объеме, например, [21]). Кроме того, в некоторых местах книги используются результаты выпуклого анализа и теории экстремальных задач. В качестве руководства по теории вероятностей рекомендуется [45], по математической статистике — [6, 7], по выпуклому анализу и теории экстремальных задач — [2, 34].

ВВЕДЕНИЕ

Эксперимент является важнейшей частью научного исследования. Мы будем рассматривать лишь математические модели эксперимента. Это означает, что реальные физические, биологические и другие эксперименты будут фигурировать в дальнейшем изложении разве лишь в виде примеров. Преимущество такого подхода состоит в общности. Каждая математическая модель оказывается приложимой во многих конкретных ситуациях.

Математические модели эксперимента тесно связаны с такими математическими дисциплинами, как теория вероятностей и математическая статистика. Отношение теории вероятностей к «действительному миру опыта» может быть определено схемой А. Н. Колмогорова [20], которая включает, в частности, задание некоторого комплекса Σ условий, допускающего неограниченное число повторений, и изучение определенного круга событий, которые могут наступать в результате осуществления этих условий. Выбор экспериментатором некоторых из условий Σ и является планированием эксперимента.

Применение теории вероятностей к изучению реальных экспериментальных ситуаций привело за последние 50 лет к замечательным научным результатам. Тем не менее возможны модели описания эксперимента, отличные от теоретико-вероятностных. Среди них важное место занимают модели, связанные с экспериментом, осуществляемым с помощью ЭВМ (имитационным экспериментом).

Математические методы планирования эксперимента оказываются во многих случаях общими для моделей, имеющих различное происхождение. Это не может вызвать удивления, если более подробно остановиться на содержании понятия «планирование эксперимента». Возможность планирования возникает в том случае, когда известно априори, что интересующий экспериментатора ответ может быть получен в результате различных, вообще говоря, экспериментов. Иными словами, имеется множество совокупностей условий $\sigma_1, \sigma_2, \dots$, причем каждая из совокупностей σ_i может дать ответ на интересующий вопрос. Предполагается, что для каждого σ_i определено значение затрат $s(\sigma_i) = s_i$, $\sigma_i \in J$ ($i \in I$), где J и I — заданные множества, и задача

состоит в выборе такого i , при котором s_i минимально. Множество J при этом может иметь достаточно сложную природу, и следует предполагать, что функция s ограничена снизу на J .

Таким образом, экспериментатор может осуществить выбор из множества I и получить нужные ему данные, создавая совокупность условий σ_i . Эксперимент, осуществляемый на основе предварительного выбора $i \in I$, принято называть *активным*. В книге будут рассматриваться лишь активные эксперименты.

Планирование активного эксперимента включает два этапа:

- 1) определение множества I и построение функций s_i ;
- 2) нахождение $i_0 \in I$, для которого s_i достигает наименьшего значения.

Последовательное решение этих двух задач и составляет содержание математической теории планирования эксперимента. Формирование критерия оптимальности в задачах планирования эксперимента имеет ряд специфических особенностей. Решение же соответствующей экстремальной задачи может часто укладываться в рамки хорошо разработанных методов. Можно тем не менее указать и специфические подходы к решению таких задач, заслуживающие отдельного рассмотрения и оказавшие влияние на общую теорию экстремальных задач и нелинейного программирования.

Для всех математических моделей мы считаем результатом эксперимента математический объект — число, множество чисел, кривую и т. п. Значительный круг прикладных задач имеет своей целью восстановление по экспериментальным данным неизвестного оператора. Предполагаются заданными два множества объектов: $\{x\} = F_1$, $\{y\} = F_2$, и оператор $A: F_1 \rightarrow F_2$. Эксперимент сопоставляет объекту x значение $\hat{y} = Ax$, которое отличается от Ax (в силу случайных экспериментальных погрешностей). Относительно F_2 разумно предположить, что для любых двух его элементов y_1 и y_2 определено расстояние $\rho(y_1, y_2)$, так что F_2 является метрическим пространством. Если эксперимент таков, что $\hat{y} \in F_2$ при каждом $x \in F_1$, то число $\rho(Ax, \hat{Ax})$ служит *мерой близости* между y и $\hat{y}$.

При восстановлении оператора A исследователь, кроме того, выбирает априори множество $\mathcal{A}$ операторов $\bar{A}$ — «моделей», определенных на F_1 , с множеством значений в F_2 . Предполагается, что в $\mathcal{A}$ можно найти оператор $\bar{A}_0$, достаточно близкий к A , и фактически строится $\bar{A}_0$.

При такой постановке задачи при каждом x рассматривают два вида погрешностей: $\rho(Ax, \hat{Ax})$ — систематическую (погрешность модели) и $\rho(\bar{A}x, \hat{Ax})$ — статистическую (случайную). При этом $\rho(Ax, \hat{Ax}) \leq \rho(Ax, \bar{A}x) + \rho(\bar{A}x, \hat{Ax})$. Эти два вида погрешностей и функция $s(x)$, определяющая стоимость проведения эксперимента при заданном значении x , обычно служат исходными для формирования критерия качества эксперимента: $K(\rho(Ax, \bar{A}x), \rho(\bar{A}x, \hat{Ax}), s(x))$. Если критерий зависит от величин и функций, которыми может распоряжаться экспериментатор,

тор (их совокупность обозначим через ξ), то естественной оказывается постановка задачи об оптимальном выборе ξ , которая и является в данном случае задачей *планирования эксперимента*. Эксперимент, для которого множество Ξ возможных значений ξ содержит по крайней мере два различных элемента, является в соответствии с принятой нами терминологией активным. Если Ξ не определено или содержит всего один элемент, то эксперимент является *пассивным*.

Относительно способов формирования критерия K можно сделать следующие общие замечания.

1. Как правило, экспериментатора интересует не один критерий, а некоторое их множество. Например, $\rho(Ax, \bar{A}x)$, $\rho(\bar{A}x, \bar{A}x)$, $s(x)$ могут быть такими критериями (считаем, что они зависят от ξ). Иными словами, более естественной является задача векторной оптимизации. Предварительный эксперимент (или теоретическое исследование) бывает необходимым для выяснения относительной значимости каждого критерия и конструирования с учетом этих сведений единого *компромиссного критерия*. Таким критерием может быть взвешенное среднее отдельных критериев, причем вес должен отражать значимость каждого из них.

2. В общем случае критерий $K(\xi)$ зависит от оператора A , подлежащего определению. Можно говорить, что задача многокритериальна — имеется столько критериев, сколько может быть различных A в совокупности $\{A\} = \mathcal{A}$. Обозначим эти критерии через $K(A, \xi)$.

Обычным способом построения компромиссного критерия является взятие верхней грани по всем возможным A , если такая верхняя грань имеет смысл, т. е. $K(\xi) = \sup_{A \in \mathcal{A}} K(A, \xi)$ (*минимаксный подход*).

Другим распространенным способом построения компромиссного критерия является осреднение критериев $K(A, \xi)$ по множеству $\mathcal{A}$. На $\mathcal{A}$ следует определить меру $\mu(dA)$ (меру предпочтительности). Если это удастся, то строится так называемый *байесовский критерий* $\bar{K}(\xi) = \int K(A, \xi) \mu(dA)$. В конечном счете приходим к необходимости решения экстремальной задачи — определению оптимального ξ , — которая может оказаться очень сложной.

Альтернативным подходом к указанным методам построения критерия может быть *последовательный подход*. Эксперимент разбивается на этапы и по мере уточнения сведений о значимости каждого критерия или о виде оператора A уточняется и сам критерий (сужается множество $\mathcal{A}$ или более четко локализуется мера μ). На каждом этапе нужно решить экстремальную задачу определения оптимального ξ для следующего этапа эксперимента.

Описанная общая схема постановки задач, связанных с планированием эксперимента по восстановлению неизвестного опе-

ратора, показывает, что эти задачи могут быть весьма сложными и требуется развитие специальных математических методов для их решения. В действительности, как мы увидим ниже, даже частные виды этой схемы связаны с содержательной теорией. Так, если искомый оператор есть функция, определенная в области из R^n (функция регрессии), известная с точностью до конечного числа неизвестных параметров, то мы приходим к классической задаче планирования регрессионного эксперимента, методы решения которой подробно излагаются в третьей главе. Другим «крайним» случаем является отсутствие случайной ошибки при наличии ошибки систематической. В частном случае, когда A является функцией, возникают задачи теории аппроксимации — обширного и развитого раздела математики.

При создании математической модели эксперимента и выборе метода планирования большую пользу может принести *имитация эксперимента с помощью ЭВМ*. Коль скоро создана некоторая модель эксперимента, безусловно, полезным является попытка его имитации. Имитация состоит в *создании алгоритма, который доставляет результаты, аналогичные результатам эксперимента*. Конечно, речь идет об эксперименте, близком к изучаемому, или эксперименте такого же типа. Так, если рассматривается задача восстановления неизвестной функции по результатам измерения ее значений, то полезно брать более или менее похожую функцию, вычислять ее значения на ЭВМ и добавлять ошибку, имитирующую процесс измерения. Ошибку генерируют с помощью датчика случайных чисел. Далее можно испытать на этой модели различные приемы планирования эксперимента и обработки данных. Для сложного эксперимента, как мы увидим ниже, теоретическое решение задачи планирования может быть очень трудным делом. Иногда эта трудность встречается даже при получении в явном виде критерия оптимальности. Здесь имитационный подход может оказаться незаменимым.

Таким образом, планирование эксперимента предполагает наличие математической модели и критерия (критериев) оптимальности. Сложный характер критериев требует специальных методов оптимизации. Ряд таких методов был специально разработан для решения задач планирования эксперимента.

Модель создается на базе априорных (или полученных в процессе предварительных экспериментов) сведений. Использование ЭВМ позволяет строить имитационные модели и отрабатывать на этих моделях методы планирования натуральных экспериментов. Наряду с этим создаются специальные методы планирования имитационного эксперимента, тесно связанные с процедурами оптимизации в статистическом моделировании.

Глава I

ЛИНЕЙНЫЙ РЕГРЕССИОННЫЙ АНАЛИЗ

Статистическая теория линейного регрессионного анализа является фундаментом многих разделов теории планирования эксперимента. В данной главе основные положения линейного регрессионного анализа подробно объясняются и обосновываются. Лицам, знакомым с теорией линейного регрессионного анализа, достаточно лишь бегло просмотреть главу с целью ознакомления с используемыми далее понятиями и обозначениями.

Материал главы изложен на языке теории матриц с использованием целого ряда результатов матричного анализа. Лицам, недостаточно знакомым со специальными фактами теории матриц, полезно перед изучением данной главы ознакомиться с § 1 приложения 1. Кроме того, читатель должен понимать основные задачи теории статистического оценивания и проверки статистических гипотез, для чего он может обратиться к [6, 7].

§ 1. Классическая линейная регрессионная модель и ее обобщения

1.1. Описательная регрессия. Одной из наиболее часто встречающихся проблем, встающих перед учеными различных специальностей, является проблема нахождения зависимости между некоторым набором величин. Эта зависимость может быть выведена из теории и (или) может быть получена на основании экспериментальных исследований. Если зависимость выведена из теоретических соображений, то довольно часто она может быть приближенно представлена в аналитическом виде, заданном с точностью до нескольких неизвестных параметров. Если же в основе построения зависимости лежат экспериментальные исследования, то параметрическая зависимость постулируется. В обоих случаях при построении математической модели должны использоваться сведения об исследуемом объекте, на основании которых мог бы быть сделан вывод о достаточной точности описания объекта моделью и, следовательно, о том, что приведенные для модели статистические выводы в определенной мере справедливы и по отношению к самому объекту.

Ниже описан «наивный» анализ зависимости внутри заданного набора наблюдаемых количественных величин.

Пусть имеется набор $m + 1$ количественных величин $y, x_1, \dots, x_m$ (например, y — характеристика получаемой продукции, x_i ($i = 1, \dots, m$) — количества поступающего сырья и характеристики производственных режимов) и N ($N \geq m$) результатов совместных измерений этих величин, которые сведены в матрицу эмпирических данных

$$(Y, F) = \begin{bmatrix} y_1 & x_{11} & \dots & x_{1m} \\ y_2 & x_{21} & \dots & x_{2m} \\ \dots & \dots & \dots & \dots \\ y_N & x_{N1} & \dots & x_{Nm} \end{bmatrix}, \quad (1.1)$$

где

$$Y = \begin{bmatrix} y_1 \\ y_2 \\ \dots \\ y_N \end{bmatrix}, \quad F = \begin{bmatrix} x_{11} & \dots & x_{1m} \\ \dots & \dots & \dots \\ x_{N1} & \dots & x_{Nm} \end{bmatrix}. \quad (1.2)$$

$y_j, x_{j1}, \dots, x_{jm}$ ($j = 1, \dots, N$) — значения величин $y, x_1, \dots, x_m$ в j -м измерении.

Представление матрицы данных в виде (1.1) означает, что величина y зависит от величин $x_1, \dots, x_m$, и основная задача анализа состоит в нахождении этой зависимости. В практических расчетах обычно $y_1, \dots, y_N$ являются результатами каких-либо измерений, а значения x_{ji} — входными для изучаемого явления, т. е. могут быть измерены до проведения измерений y_j .

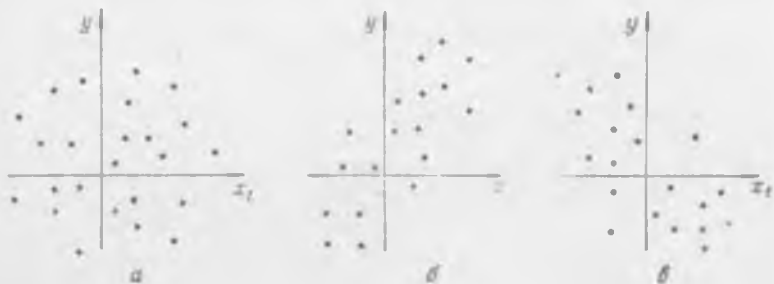


Рис. 1

или даже могут быть выбраны экспериментатором. Этим обусловлено название вектора Y и матрицы F : Y называется *вектором результатов измерений*, а матрица F — *матрицей плана эксперимента* или просто *матрицей плана*.

Если имеются только две переменные y и x_1 (т. е. $m = 1$), то эмпирические данные могут быть представлены графически (рис. 1). Подбор прямой линии к набору измерений — одна из наиболее часто встречающихся процедур, применяемых во многих отраслях науки; графики, подобные приведенному на рис. 1,

являются распространенным средством определения статистической зависимости.

Если y и x_1 трактовать как случайные величины, то на рис. 1 в случае *б* имеет место положительная корреляция между y и x_1 , в случае *в* — отрицательная, в случае *а* корреляция отсутствует. Напомним, что наличие у случайных величин ξ и η положительной (отрицательной) корреляции означает, что в среднем, чем больше значения ξ , тем больше (соответственно меньше) значения η .

В подавляющем большинстве случаев может быть найдена лишь достаточно сложная функция (например, полином высокой степени), проходящая через каждую точку графика, но на практике обычно требуется найти относительно простую функцию, которая с удовлетворительной точностью аппроксимирует предполагаемую статистическую зависимость.

В общей постановке задача описания эмпирической зависимости с помощью параметрической регрессии предполагает, что задается функция, определенная с точностью до нескольких параметров, которые подбирают таким образом, чтобы получающаяся функция с максимальной точностью соответствовала матрице данных (1.1). Функция η при этом называется *эмпирической регрессией*. Если η линейна как функция неизвестных параметров, то регрессия называется *линейной* (в противном случае — *нелинейной*).

Простейшим примером линейной регрессии является

$$\eta(x) = \eta(x_1, \dots, x_m, \theta_1, \dots, \theta_m) = \theta_1 x_1 + \dots + \theta_m x_m, \quad (1.3)$$

где $\theta_1, \dots, \theta_m$ — неизвестные параметры регрессии. В частном случае, когда значения величины x_1 тождественно равны единице, функция регрессии (1.3) принимает вид $\eta(x) = \theta_1 + \theta_2 x + \dots + \theta_m x_m$ или — после очевидной замены нумерации — вид $\eta(x) = \theta_0 + \theta_1 x_1 + \dots + \theta_k x_k$. Если линейность кажется неподходящей, часто можно найти преобразования над данными, которые приводят к приблизительной линейности (подобные преобразования прекрасно описаны в [30]), и поэтому линейные зависимости для практики более важны, чем это может показаться на первый взгляд.

Следующей проблемой после выбора типа функции η является определение таких численных значений неизвестных параметров $\theta_1, \dots, \theta_m$, при которых функция регрессии будет достаточно хорошо (или даже наилучшим образом) описывать эмпирические данные. Для ее решения прежде всего необходимо задать критерий, который определял бы степень соответствия эмпирических данных и регрессионной зависимости. Любой такой критерий должен учитывать отклонения между измеренными значениями $y_1, \dots, y_N$ и приближенными значениями

$$\theta_1 = \sum_{i=1}^m x_i \theta_i. \quad (1.4)$$

Следовательно, необходимо учитывать отклонения ε_j ($j = 1, \dots, N$), определяемые по формуле

$$\varepsilon_j = y_j - \hat{y}_j, \quad j = 1, \dots, N. \quad (1.5)$$

С учетом (1.4) формула (1.5) переписывается в виде

$$y_j = \sum_{i=1}^m x_{ji} \theta_i + \varepsilon_j, \quad j = 1, \dots, N, \quad (1.6)$$

В матричной форме это записывается так:

$$Y = F\theta + \varepsilon, \quad (1.7)$$

где $\theta = (\theta_1, \dots, \theta_m)^T$ — вектор неизвестных параметров, $\varepsilon = (\varepsilon_1, \dots, \varepsilon_N)^T$ — вектор отклонений. При заданной матрице (1.1) вектор ε зависит от значений θ , и поэтому от удачного выбора θ зависит качество регрессионного описания. Критерии оптимального выбора θ могут быть заданы различными способами. Например, наилучшим могут считаться те θ , для которых величины

$$\max_{j=1, \dots, N} |\varepsilon_j|, \quad \sum_{j=1}^N |\varepsilon_j|, \quad \sum_{j=1}^N \varepsilon_j^2 = \varepsilon^T \varepsilon$$

минимальны. Ниже рассматривается только третий из указанных критериев. Это объясняется тем, что получающиеся формулы расчета значений θ относительно просты с вычислительной точки зрения, а сами эти значения обладают определенными свойствами оптимальности и хорошо зарекомендовали себя на практике.

По определению вектор

$$\hat{\theta} = \arg \min_{\theta \in \mathbb{R}^m} (Y - F\theta)^T (Y - F\theta) = \arg \min_{\theta \in \mathbb{R}^m} \varepsilon^T \varepsilon \quad (1.8)$$

называется (эмпирической) оценкой метода наименьших квадратов (МНК).

Здесь впервые встретился символ $\arg \min$, который неоднократно будет использоваться в дальнейшем и обозначает следующее. Пусть на множестве X задана ограниченная снизу функция $f(x)$, которая на этом множестве достигает минимального значения. Тогда под $\arg \min_{x \in X} f(x)$ понимается любая точка

минимума функции f на множестве X , т. е. такая точка x_* , которая принадлежит множеству X и в которой выполнено $f(x_*) = \min_{x \in X} f(x)$.

Одно из важнейших свойств оценки МНК сформулировано ниже и заключается в том, что задача ее вычисления сводится к задаче решения системы линейных алгебраических уравнений $F^T F \hat{\theta} = F^T Y$, которая называется *системой нормальных уравнений*.

Лемма 1.1. Вектор (1.8) является решением системы уравнений

$$F^T F \theta = F^T Y. \quad (1.9)$$

Доказательство. Примем обозначение

$$Q(\theta) = (Y - F\theta)^T (Y - F\theta) = \sum_{j=1}^N \left(y_j - \sum_{i=1}^m x_{ji} \theta_i \right)^2.$$

Функция $Q(\theta)$ задана на R^m и является неотрицательно определенной квадратичной формой. Поэтому у нее существует конечный минимум $Q_{\min}$, и, чтобы получить систему уравнений для точки минимума функции $Q(\theta)$ (этой точкой является (1.8)), достаточно продифференцировать $Q(\theta)$ по $\theta_1, \dots, \theta_m$ и полученные производные приравнять нулю.

Для вычисления производных можно воспользоваться результатами из § 3 приложения 1, но мы вычислим эти производные непосредственно.

Имеем

$$\begin{aligned} \frac{\partial Q(\theta)}{\partial \theta_i} &= \partial \left[\sum_{j=1}^N \left(y_j - \sum_{i=1}^m x_{ji} \theta_i \right)^2 \right] / \partial \theta_i = \\ &= \sum_{j=1}^N 2 \left(y_j - \sum_{i=1}^m x_{ji} \theta_i \right) \left[\partial \left(y_j - \sum_{i=1}^m x_{ji} \theta_i \right) / \partial \theta_i \right] = \\ &= 2 \sum_{j=1}^N \left(y_j - \sum_{i=1}^m x_{ji} \theta_i \right) (-x_{ji}) = 2 \left[\sum_{j=1}^N \sum_{i=1}^m x_{ji} x_{ji} \theta_i - \sum_{j=1}^N x_{ji} y_j \right]. \end{aligned}$$

Записывая полученное выражение в матричном виде, получаем

$$\nabla Q(\theta) = \left(\frac{\partial Q(\theta)}{\partial \theta_1}, \dots, \frac{\partial Q(\theta)}{\partial \theta_m} \right)^T = 2 [F^T F \theta - F^T Y].$$

Приравнявая $\nabla Q(\theta)$ нулю, получаем (1.9). Лемма доказана.

Поскольку число неизвестных в системе нормальных уравнений (1.9) не меньше, чем число линейно независимых уравнений, эта система всегда имеет по крайней мере одно решение.

Если матрица $F^T F$ невырождена, то система нормальных уравнений (1.9) имеет единственное решение

$$\hat{\theta} = (F^T F)^{-1} F^T Y, \quad (1.10)$$

которое однозначно определяет оценку МНК.

Пример 1.1. Предположим, что результаты измерений переменной y в точке x равны $y(x) = x^2$, измерены значения y в точках $x_{(1)} = 0$, $x_{(2)} = 1/2$, $x_{(3)} = 1$, а аппроксимируется кривая $y(x)$ прямой $\eta(x) = \theta_1 + \theta_2 x$. В данном примере переменная x_1 тождественно равна единице, переменная x_2 заменена на x ,

$$Y = \begin{bmatrix} 0 \\ 1/4 \\ 1 \end{bmatrix}, \quad F = \begin{bmatrix} 1 & 0 \\ 1 & 1/2 \\ 1 & 1 \end{bmatrix}.$$

Вычислим МНК-оценку (1.10). Имеем

$$F^T F = \begin{bmatrix} 1 & 1 & 1 \\ 0 & 1/2 & 1 \end{bmatrix} \begin{bmatrix} 1 & 0 \\ 1 & 1/2 \\ 1 & 1 \end{bmatrix} = \begin{bmatrix} 3 & 3/2 \\ 3/2 & 5/4 \end{bmatrix},$$

$$\det F^T F = \frac{3}{2}, \quad (F^T F)^{-1} = \begin{bmatrix} \frac{5}{4} \cdot \frac{2}{3} & -\frac{3}{2} \cdot \frac{2}{3} \\ -\frac{3}{2} \cdot \frac{2}{3} & 3 \cdot \frac{2}{3} \end{bmatrix} = \begin{bmatrix} 5/6 & -1 \\ -1 & 2 \end{bmatrix},$$

$$F^T Y = \begin{bmatrix} 1 & 1 & 1 \\ 0 & 1/2 & 1 \end{bmatrix} \begin{bmatrix} 0 \\ 1/4 \\ 1 \end{bmatrix} = \begin{bmatrix} 5/4 \\ 9/8 \end{bmatrix},$$

$$\hat{\theta} = (F^T F)^{-1} F^T Y = \begin{bmatrix} 5/6 & -1 \\ -1 & 2 \end{bmatrix} \begin{bmatrix} 5/4 \\ 9/8 \end{bmatrix} = \begin{bmatrix} -1/12 \\ 1 \end{bmatrix}.$$

Таким образом, по результатам вычисления функции $y(x) = x^2$ в точках 0, 1/2, 1 указанная функция приближается линейной функцией $(-1/12 + x)$, рассчитанной по МНК (рис. 2).

Математической моделью рассмотренной выше схемы описательной регрессии является классическая линейная регрессионная модель, определение которой дается в п. 1.2 и изучение которой составляет суть данного параграфа.

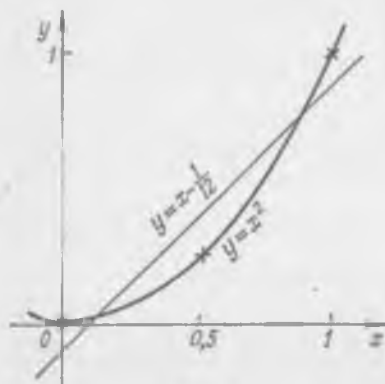


Рис. 2

1.2. Основные определения и некоторые свойства оценок. В матричном виде произвольная линейная регрессионная модель (схема линейной регрессии) представляется как

$$Y = F\theta + \varepsilon, \quad \varepsilon \sim (0, \Sigma), \quad (1.11)$$

где $Y = (y_1, \dots, y_N)^T$ — случайный вектор результатов измерений (N — число измерений); $F = \|x_{ij}\|$ — фиксированная $N \times m$ -

матрица плана (вектор Y и матрица F имеют вид (1.2)); $\theta = (\theta_1, \dots, \theta_m)^T$ — вектор неизвестных параметров, которые должны быть оценены по результатам измерений $y_1, \dots, y_N$; $\varepsilon = (\varepsilon_1, \dots, \varepsilon_N)^T$ — вектор ненаблюдаемых случайных ошибок*), математические ожидания которых равны нулю ($E\varepsilon = 0$); Σ — фиксированная, положительно определенная $N \times N$ -матрица, являющаяся дисперсионной матрицей случайных ошибок ($\Sigma = E\varepsilon\varepsilon^T = \|E\varepsilon_i\varepsilon_j\|_{i,j=1}^N$).

Обозначение $\varepsilon \sim (0, \Sigma)$ — сокращенная запись того, что случайный вектор ε имеет распределение с нулевым вектором сред-

*) Т.е. случайных величин, определенных на одном и том же вероятностном пространстве.

них ($E\varepsilon = 0$) и дисперсионной (ковариационной) матрицей $E\varepsilon\varepsilon^T = \Sigma$. При необходимости проверки гипотез знания первых двух моментов случайного вектора ε недостаточно: необходимо полностью знать его распределение (см. п. 1.6 и § 3).

Удобное краткое обозначение модели (1.11), которое далее используется, — упорядоченная тройка $(Y, F\theta, \Sigma)$.

Отличие модели $(Y, F\theta, \Sigma)$ от рассмотренной выше модели описательной регрессии (1.7) состоит в дополнительных предположениях о случайности вектора отклонений ε (который теперь называется *вектором случайных ошибок*) и о том, что $E\varepsilon = 0$, $D\varepsilon = \Sigma$. Эти предположения будут ниже служить основой для изучения свойств оценок МНК.

Так же как и для модели описательной регрессии (1.7), при использовании схемы $(Y, F\theta, \Sigma)$ подробная запись результатов измерений $y_1, \dots, y_N$ имеет вид (1.6). Это соответствует тому, что функция регрессии *) имеет вид (1.3).

Далее неоднократно будут исследоваться линейные регрессионные модели, описанные в следующих трех примерах.

Пример 1.2. Линейная регрессия на отрезке:

$$\eta(x) = \theta_1 + \theta_2 x, \quad x \in [-1, 1]. \quad (1.12)$$

В этом случае модель (1.6) для результатов измерений записывается в виде

$$y_j = \theta_1 + \theta_2 x_j + \varepsilon_j, \quad j = 1, \dots, N, \quad (1.13)$$

или в матричном виде: $Y = F\theta + \varepsilon$, где

$$Y = \begin{bmatrix} y_1 \\ \vdots \\ y_N \end{bmatrix}, \quad F = \begin{bmatrix} 1 & x_1 \\ \vdots & \vdots \\ 1 & x_N \end{bmatrix}, \quad \varepsilon = \begin{bmatrix} \varepsilon_1 \\ \vdots \\ \varepsilon_N \end{bmatrix}, \quad \theta = \begin{bmatrix} \theta_1 \\ \theta_2 \end{bmatrix}. \quad (1.14)$$

В данном случае величины $x_1, \dots, x_N$ представляют собой значения переменной x , которая совпадает с переменной x_2 из формулы (1.3). Переменная x_1 из формулы (1.3) принимает значения, тождественно равные единице.

Пример 1.3. Квадратичная регрессия на отрезке:

$$\eta(x) = \theta_1 + \theta_2 x + \theta_3 x^2, \quad x \in [-1, 1]. \quad (1.15)$$

Для матричной модели результатов измерений $Y = F\theta + \varepsilon$ векторы Y и ε те же, что и в (1.14),

$$\theta = (\theta_1, \theta_2, \theta_3)^T, \quad F = \begin{bmatrix} 1 & x_1 & x_1^2 \\ \vdots & \vdots & \vdots \\ 1 & x_N & x_N^2 \end{bmatrix}. \quad (1.16)$$

*) Функцией регрессии называется функция, вычисляемая со случайной ошибкой, математическое ожидание которой равно нулю.

Если представить функцию регрессии (1.15) в виде (1.3), то в качестве переменной x_1 следует взять постоянную функцию, тождественно равную единице, в качестве x_2 — функцию x , а в качестве x_3 — функцию x^2 .

Пример 1.4. Сформулируем схему регрессии для задачи определения весов $\theta_1, \theta_2, \theta_3$ трех предметов на весах с двумя чашками. Предположим, что результат j -го измерения есть разность веса содержимого второй и первой чашек плюс некоррелированная с остальными случайная ошибка e_j со средним 0 и дисперсией σ^2 , т. е.

$$y_j = \theta_1 x_{1j} + \theta_2 x_{2j} + \theta_3 x_{3j} + e_j, \\ E e_j = 0, \quad E e_j^2 = \sigma^2, \quad E e_i e_j = 0, \quad i \neq j,$$

где

$$x_{lj} = \begin{cases} (-1)^{m_{lj}}, & \text{если } l\text{-й груз взвешивался в } j\text{-м эксперименте,} \\ 0 & \text{в противном случае,} \end{cases}$$

m_{lj} — номер чашки, на которой в j -м взвешивании лежал l -й груз ($l = 1, 2, 3$).

В данном случае функция регрессии имеет вид

$$\eta(x) = \theta_1 x_1 + \theta_2 x_2 + \theta_3 x_3, \quad (1.17)$$

где переменные x_1, x_2 и x_3 могут принимать лишь три значения: $-1, 0, 1$.

Ниже часто будет использоваться класс линейных однородных оценок, т. е. оценок $\bar{\theta}$ вида $\bar{\theta} = CY$, где C — произвольная матрица порядка $m \times N$. При этом слово «однородная» будет опускаться. Линейность оценки $\bar{\theta}$ означает, что все компоненты этой оценки — линейные комбинации результатов измерений $y_1, \dots, y_N$.

Одним из важнейших свойств оценок является условие несмещенности*). Оказывается, что необходимое и достаточное условие несмещенности линейных оценок параметров θ имеет достаточно простой вид.

Лемма 1.2. Для схемы линейной регрессии $(Y, F\theta, \Sigma)$ линейная оценка $\bar{\theta} = CY$ является несмещенной оценкой параметров θ , если и только если

$$CF = I_m. \quad (1.18)$$

Доказательство. Поскольку $EY = F\theta$, то $E\bar{\theta} = E CY = CF\theta$. Покажем, что выполнение

$$CF\theta = \theta \text{ для всех } \theta \in R^m \quad (1.19)$$

равносильно выполнению (1.18). Действительно, умножая справа обе части матричного равенства (1.18) на θ , получаем (1.19).

*) Оценка $\bar{\theta}$ параметров θ называется несмещенной, если $E\bar{\theta} = \theta$ для всех θ .

С другой стороны, взяв в (1.19) в качестве θ вектор $c_i = (0, \dots, 0, 1, 0, \dots, 0)$, все компоненты которого равны нулю, кроме i -го, равного единице, получим векторное равенство, являющееся i -м столбцом матричного равенства (1.18). Проделав это для всех $i = 1, \dots, m$, получим (1.18). Лемма доказана.

Обозначим: $\Xi(\theta)$ — множество всевозможных линейных несмещенных оценок параметров θ ; $D\bar{\theta}$ — дисперсионная матрица оценки $\bar{\theta}$. Наилучшей линейной несмещенной (НЛН) оценкой параметров θ называется (в случае ее существования) такая оценка $\hat{\theta} \in \Xi(\theta)$, что $D\hat{\theta} \leq D\bar{\theta}^*$ для всех $\bar{\theta} \in \Xi(\theta)$.

Ниже будет показано, что для схемы регрессии $(Y, F\theta, \Sigma)$ НЛН-оценка существует и совпадает с оценкой, полученной по методу наименьших квадратов (МНК-оценкой).

Из свойств неотрицательно определенных матриц следует, что если $\bar{\theta} = (\bar{\theta}_1, \dots, \bar{\theta}_m)$ есть НЛН-оценка параметров $\theta = (\theta_1, \dots, \theta_m)^T \in \mathbb{R}^m$, то $\det D\bar{\theta} \leq \det D\hat{\theta}$, $a^T [D\hat{\theta}] a \leq a^T [D\bar{\theta}] a$ и $[D\hat{\theta}]_{ii} \leq [D\bar{\theta}]_{ii}$ для всех $\bar{\theta} \in \Xi(\theta)$, $a \in \mathbb{R}^m$ ($i = 1, \dots, m$) (здесь $[D]_{ii}$ — диагональные элементы матрицы D). Отсюда следует, в частности, что оценка $\hat{\theta}_i$ каждой компоненты θ_i вектора θ обладает наименьшей дисперсией (в классе $\Xi(\theta)$) и для любого $a \in \mathbb{R}^m$ оценка $a^T \hat{\theta}$ является НЛН-оценкой скалярного параметра $a^T \theta$.

Пусть задана линейная регрессионная модель $(Y, F\theta, \Sigma)$, где матрица F имеет полный ранг m , т. е. $\text{rang } F = m$ (при этом говорят, что модель невырождена или является моделью полного ранга). Стандартной МНК-оценкой параметров θ называется вектор

$$\hat{\theta} = (F^T F)^{-1} F^T Y, \quad (1.20)$$

который по виду совпадает с эмпирической оценкой МНК (1.10). Слово «стандартная» обычно будет опускаться. Своим названием оценка (1.20) обязана тому, что в силу леммы 1.1

$$\hat{\theta} = \arg \min_{\theta \in \mathbb{R}^m} (Y - F\theta)^T (Y - F\theta). \quad (1.21)$$

Очевидно, что оценка (1.20) является линейной. Кроме того, легко видеть, что эта оценка несмещенная. Действительно, для всех θ имеем

$$E\hat{\theta} = (F^T F)^{-1} F^T (EY) = (F^T F)^{-1} F^T F\theta = \theta,$$

что эквивалентно условию несмещенности оценки (1.20). Приведем примеры вычисления стандартных МНК-оценок.

* Неравенство $A \geq B$ ($A > B$) для квадратных матриц A, B одного порядка означает, что матрица $(A - B)$ неотрицательно определена (соответственно положительно определена).

Пример 1.5 (МНК-оценка для линейной регрессии на отрезке). Рассмотрим линейную регрессионную модель $(Y, F\theta, \sigma^2 I_N)$, где Y, F, θ определяются по формулам (1.14). Выведем явный вид МНК-оценки $\hat{\theta} = (\hat{\theta}_1, \hat{\theta}_2)^T$. Имеем

$$F^T F = \begin{bmatrix} 1 & 1 & \dots & 1 \\ x_1 & x_2 & \dots & x_N \end{bmatrix} \begin{bmatrix} 1 & \dots & 1 \\ x_1 & \dots & x_N \end{bmatrix} = \begin{bmatrix} N & \sum_{i=1}^N x_i \\ \sum_{i=1}^N x_i & \sum_{i=1}^N x_i^2 \end{bmatrix},$$

$$\det F^T F = N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2,$$

$$(F^T F)^{-1} = (\det F^T F)^{-1} \begin{bmatrix} \sum_{i=1}^N x_i^2 & -\sum_{i=1}^N x_i \\ -\sum_{i=1}^N x_i & N \end{bmatrix},$$

$$F^T Y = \begin{bmatrix} 1 & \dots & 1 \\ x_1 & \dots & x_N \end{bmatrix} \begin{bmatrix} y_1 \\ \dots \\ y_N \end{bmatrix} = \begin{bmatrix} \sum_{i=1}^N y_i \\ \sum_{i=1}^N x_i y_i \end{bmatrix},$$

$$\begin{aligned} \hat{\theta} &= (F^T F)^{-1} F^T Y = (\det F^T F)^{-1} \begin{bmatrix} \sum_{i=1}^N x_i^2 & -\sum_{i=1}^N x_i \\ -\sum_{i=1}^N x_i & N \end{bmatrix} \begin{bmatrix} \sum_{i=1}^N y_i \\ \sum_{i=1}^N x_i y_i \end{bmatrix} = \\ &= (\det F^T F)^{-1} \begin{bmatrix} \sum_{i=1}^N x_i^2 \sum_{i=1}^N y_i - \sum_{i=1}^N x_i \sum_{i=1}^N x_i y_i \\ N \sum_{i=1}^N x_i y_i - \sum_{i=1}^N x_i \sum_{i=1}^N y_i \end{bmatrix}. \end{aligned}$$

Следовательно, в общем виде МНК-оценки $\hat{\theta}_1, \hat{\theta}_2$ параметров θ_1, θ_2 линейной регрессии $(\theta_1 + \theta_2 x)$ вычисляются по формулам

$$\hat{\theta}_1 = (\sum_{i=1}^N x_i^2 \sum_{i=1}^N y_i - \sum_{i=1}^N x_i \sum_{i=1}^N x_i y_i) / [N \sum_{i=1}^N x_i^2 - (\sum_{i=1}^N x_i)^2], \quad (1.22)$$

$$\hat{\theta}_2 = (N \sum_{i=1}^N x_i y_i - \sum_{i=1}^N x_i \sum_{i=1}^N y_i) / [N \sum_{i=1}^N x_i^2 - (\sum_{i=1}^N x_i)^2], \quad (1.23)$$

где все суммы берутся по j от 1 до N .

Обычно рассматривается такой способ выбора точек проведения измерений x_j ($j = 1, \dots, N$), для которого выполнено равенство

$$\sum_{i=1}^N x_i = 0. \quad (1.24)$$

Это условие выполнено, например, если измерения проводятся в точках, расположенных симметрично относительно нуля. При

выполнении условия (1.24) формулы (1.22) и (1.23) упрощаются и принимают вид

$$\hat{\theta}_1 = \frac{1}{N} \sum_{j=1}^N y_j, \quad (1.25)$$

$$\hat{\theta}_2 = \frac{\sum_{j=1}^N x_j y_j}{\sum_{j=1}^N x_j^2}. \quad (1.26)$$

Формула (1.25) означает, что при выполнении (1.24) оценка $\hat{\theta}_1$ для θ_1 является средним арифметическим результатов измерений.

Вычислим теперь оценки $\hat{\theta}_1, \hat{\theta}_2$ параметров регрессии $\theta_1 + \theta_2 x$ в случае, когда $N=3$, а измерения проводятся в трех точках: $x_1 = -1, x_2 = 0, x_3 = 1$. Имеем

$$F = \begin{bmatrix} 1 & -1 \\ 1 & 0 \\ 1 & 1 \end{bmatrix}, \quad F^T = \begin{bmatrix} 1 & 1 & 1 \\ -1 & 0 & 1 \end{bmatrix}, \quad F^T F = \begin{bmatrix} 3 & 0 \\ 0 & 2 \end{bmatrix},$$

$$(F^T F)^{-1} = \begin{bmatrix} 1/3 & 0 \\ 0 & 1/2 \end{bmatrix}, \quad F^T Y = \begin{bmatrix} y_1 + y_2 + y_3 \\ y_3 - y_1 \end{bmatrix},$$

$$\hat{\theta} = \begin{bmatrix} (y_1 + y_2 + y_3)/3 \\ (y_3 - y_1)/2 \end{bmatrix}.$$

Конечно, эта формула — частный случай формул (1.25), (1.26).

Пример 1.6 (МНК-оценка для квадратичной регрессии). Рассмотрим функцию регрессии (1.15) и предположим, что измерения проводятся в точках x_j , симметрично расположенных относительно нуля. Для этих точек будут выполнены соотношения

$$\sum_{j=1}^N x_j = 0, \quad \sum_{j=1}^N x_j^3 = 0, \quad (1.27)$$

одно из которых совпадает с (1.24). Вычислим МНК-оценку параметров $\theta = (\theta_1, \theta_2, \theta_3)^T$.

Используя (1.16) и (1.27), имеем

$$F^T F = \begin{bmatrix} 1 & \dots & 1 \\ x_1 & \dots & x_N \\ x_1^2 & \dots & x_N^2 \end{bmatrix} \begin{bmatrix} 1 & x_1 & x_1^2 \\ \vdots & \vdots & \vdots \\ 1 & x_N & x_N^2 \end{bmatrix} =$$

$$= \begin{bmatrix} N & \sum x_j & \sum x_j^2 \\ \sum x_j & \sum x_j^2 & \sum x_j^3 \\ \sum x_j^2 & \sum x_j^3 & \sum x_j^4 \end{bmatrix} = \begin{bmatrix} N & 0 & \sum x_j^2 \\ 0 & \sum x_j^2 & 0 \\ -\sum x_j^2 & 0 & \sum x_j^4 \end{bmatrix}.$$

Примем

$$a = \sum_{i=1}^N x_i^2, \quad b = \sum_{i=1}^N x_i^4, \quad c = Nab - a^3 \quad (1.28)$$

и предположим, что $c = \det F^T F \neq 0$. Тогда

$$(F^T F)^{-1} = \begin{bmatrix} ab/c & 0 & -a^2/c \\ 0 & Nb - a^2/c & 0 \\ -a^2/c & 0 & Na/c \end{bmatrix}, \quad (1.29)$$

$$F^T Y = \begin{bmatrix} \sum y_i \\ \sum x_i y_i \\ \sum x_i^2 y_i \end{bmatrix},$$

$$\hat{\theta} = \begin{bmatrix} \hat{\theta}_1 \\ \hat{\theta}_2 \\ \hat{\theta}_3 \end{bmatrix} = \begin{bmatrix} \frac{a}{c} [b \sum y_i - a \sum x_i^2 y_i] \\ \frac{Nb - a^2}{c} \sum x_i y_i \\ \frac{a}{c} [N \sum x_i^2 y_i - a \sum y_i] \end{bmatrix}. \quad (1.30)$$

В частности, если $N = 3$ и измерения проведены в точках $x_1 = -1$, $x_2 = 0$, $x_3 = 1$, то из (1.30) получаем

$$\begin{aligned} \hat{\theta}_1 &= y_2, \\ \hat{\theta}_2 &= \frac{1}{2} (y_3 - y_1), \\ \hat{\theta}_3 &= \frac{1}{2} [3(y_1 + y_3) - 2(y_1 + y_2 + y_3)] = \frac{1}{2} (y_1 + y_3) - y_2. \end{aligned} \quad (1.31)$$

Впрочем, формулы (1.31) легко получить и непосредственно.

Тип дисперсионной матрицы $DY = \Sigma$ определяет как название линейной регрессионной модели $(Y, F\theta, \Sigma)$, так и оптимальные способы оценивания неизвестных параметров. Наиболее простой случай — это случай, когда ошибки $\varepsilon = (\varepsilon_1, \dots, \varepsilon_N)^T$ некоррелированы и имеют одинаковую дисперсию, т. е. $\Sigma = \sigma^2 I_N$, где $\sigma^2 > 0$ — дисперсия одного измерения (возможно, неизвестная). В этом случае модель $(Y, F\theta, \sigma^2 I_N)$ называется *классической линейной регрессионной моделью*.

В п. 1.3 доказан один из основных теоретических фактов линейного регрессионного анализа, согласно которому для классической линейной регрессионной модели НЛН-оценкой является стандартная МНК-оценка.

1.3. Теорема Гаусса — Маркова для классической линейной регрессии. Прежде чем доказывать основную теорему, покажем, как вычисляются дисперсионные матрицы линейных несмещенных оценок.

Лемма 1.3. Пусть ξ — некоторый случайный вектор размерности p , $D\xi = E(\xi - E\xi)(\xi - E\xi)^T$ — его конечная дисперсионная матрица, a — произвольный детерминированный q -вектор и L — фиксированная $q \times p$ -матрица. Тогда дисперсионная матрица случайного вектора $\eta = L\xi + a$ равна $D\eta = LD\xi L^T$.

Доказательство. По определению дисперсионной матрицы имеем

$$\begin{aligned} D\eta &= E(\eta - E\eta)(\eta - E\eta)^T = E(L\xi + a - EL\xi - a) \times \\ &\quad \times (L\xi + a - EL\xi - a)^T = EL(\xi - E\xi)[L(\xi - E\xi)]^T = \\ &= L[E(\xi - E\xi)(\xi - E\xi)^T]L^T = LD\xi L^T. \end{aligned}$$

Лемма доказана.

Лемма 1.4. Пусть задана классическая линейная регрессионная модель $(Y, F\theta, \sigma^2 I_N)$ и $\bar{\theta} = CY$ — произвольная линейная статистика (оценка). Тогда

$$D\bar{\theta} = \sigma^2 CC^T. \quad (1.32)$$

Доказательство. Используя лемму 1.3 с заменой ξ на e , a на $F\theta$, L на C , имеем

$$D\bar{\theta} = DCY = C[DY]C^T = \sigma^2 CI_N C^T = \sigma^2 CC^T.$$

Лемма доказана.

Из утверждения леммы 1.4 следует, что если задана классическая линейная регрессионная модель полного ранга, то дисперсионная матрица стандартной МНК-оценки (1.20) равна

$$D\hat{\theta} = \sigma^2 (F^T F)^{-1}. \quad (1.33)$$

Действительно, $\hat{\theta} = C_0 Y$, где $C_0 = (F^T F)^{-1} F^T$; поэтому из (1.32) получаем

$$\begin{aligned} D\hat{\theta} &= \sigma^2 C_0 C_0^T = \sigma^2 (F^T F)^{-1} F^T [(F^T F)^{-1} F^T]^T = \\ &= \sigma^2 (F^T F)^{-1} F^T F (F^T F)^{-1} = \sigma^2 (F^T F)^{-1}. \end{aligned}$$

Теорема 1.1 (теорема Гаусса — Маркова). Пусть задана классическая линейная регрессионная модель $(Y, F\theta, \sigma^2 I_N)$, где $N \times m$ -матрица F имеет полный ранг, равный m . Тогда МНК-оценка (1.20) является НЛН-оценкой.

Доказательство. Пусть $\bar{\theta} = CY$ — произвольная линейная несмещенная оценка параметров θ , а $\hat{\theta} = C_0 Y$ — МНК-оценка. Здесь $C_0 = (F^T F)^{-1} F^T$. В силу леммы 1.2 условия несмещенности оценок θ и $\hat{\theta}$ записываются в виде

$$CF = I_m, \quad CF = I_m. \quad (1.34)$$

Положим $A = C - C_0$. Из (1.34) получаем

$$AF = (C - C_0)F = CF - C_0 F = I_m - I_m = 0,$$

откуда $AC_0^T = AF(F^T F)^{-1} = 0$. Используя полученное равенство и выражения (1.32), (1.33) для дисперсионных матриц оценок $\bar{\theta}$ и $\hat{\theta}$, получаем

$$\begin{aligned} D\bar{\theta} &= \sigma^2 CC^T = \sigma^2 [C_0 + A][C_0^T + A^T] = \sigma^2 [C_0 C_0^T + C_0 A^T + AC_0^T + \\ &\quad + AA^T] = \sigma^2 C_0 C_0^T + [AC_0^T + (AC_0^T)^T] + \sigma^2 AA^T = D\hat{\theta} + \sigma^2 AA^T. \end{aligned}$$

Так как матрица AA^T неотрицательно определена, то $D\bar{\theta} \geq D\hat{\theta}$. Теорема доказана.

1.4. Оценивание дисперсии в классической линейной регрессионной модели. В данном пункте приведена оценка неизвестной дисперсии одного измерения σ^2 в невырожденной классической линейной регрессионной модели $(Y, F\theta, \sigma^2 I_N)$.

Теорема 1.2. Несмещенной оценкой для σ^2 в невырожденной модели $(Y, F\theta, \sigma^2 I_N)$ является статистика

$$s^2 = (Y - F\hat{\theta})^T (Y - F\hat{\theta}) / (N - m), \quad (1.35)$$

где $\hat{\theta}$ есть МНК-оценка (1.20).

Доказательство. Положим $\hat{e} = Y - F\hat{\theta}$. Имеем

$$\hat{e} = Y - F(F^T F)^{-1} F^T Y = F\theta + \varepsilon - F(F^T F)^{-1} F^T (F\theta + \varepsilon) =$$

$$= F\theta + \varepsilon - F(F^T F)^{-1} F^T F\theta - F(F^T F)^{-1} F^T \varepsilon = \varepsilon - F(F^T F)^{-1} F^T \varepsilon,$$

т. е. $\hat{e} = G\varepsilon$, где $G = I_N - F(F^T F)^{-1} F^T$. Легко проверяется, что $G = G^T$ и $G^2 = G$ (т. е. матрица G идемпотентна). Отсюда получаем

$$\hat{e}^T \hat{e} = \varepsilon^T G^T G \varepsilon = \varepsilon^T G G \varepsilon = \varepsilon^T G \varepsilon.$$

Согласно теореме 1.10 из приложения 1, для любых согласованных прямоугольных матриц A, B справедливо равенство

$$\text{tr } AB = \text{tr } BA. \quad (1.36)$$

Применяя (1.36) с $A = \varepsilon^T, B = G\varepsilon$ и учитывая, что операции взятия математического ожидания и следа перестановочны (по определению интеграл от матрицы — матрица из интегралов см. с. 308), получаем

$$\begin{aligned} \mathbf{E} \hat{e}^T \hat{e} &= \mathbf{E} \text{tr } \varepsilon^T G \varepsilon = \mathbf{E} \text{tr } G \varepsilon \varepsilon^T = \text{tr } G (\mathbf{E} \varepsilon \varepsilon^T) = \\ &= \sigma^2 \text{tr } G = \sigma^2 [\text{tr } I_N - \text{tr } F(F^T F)^{-1} F^T] \end{aligned}$$

Снова применяя (1.36) (теперь с $A = F, B = (F^T F)^{-1} F^T$), получаем

$$\mathbf{E} \hat{e}^T \hat{e} = \sigma^2 [\text{tr } I_N - \text{tr } I_m] = \sigma^2 (N - m). \quad (1.37)$$

Это соотношение эквивалентно тому, что $\mathbf{E} s^2 = \sigma^2$ (т. е. несмещенности оценки s^2). Теорема доказана.

Пример 1.7 (продолжение примера 1.5). Пусть в обозначениях примера 1.5 справедливо соотношение (1.24). Тогда оценки неизвестных параметров $\theta = (\theta_1, \theta_2)^T$ регрессии $\theta_1 + \theta_2 x$ определяются по формулам (1.25), (1.26). Далее

$$\begin{aligned} Y - F\hat{\theta} &= \begin{bmatrix} y_1 \\ \vdots \\ y_N \end{bmatrix} - \begin{bmatrix} 1 & x_1 \\ \vdots & \vdots \\ 1 & x_N \end{bmatrix} \begin{bmatrix} \hat{\theta}_1 \\ \hat{\theta}_2 \end{bmatrix} = \begin{bmatrix} y_1 - \hat{\theta}_1 - \hat{\theta}_2 x_1 \\ \vdots \\ y_N - \hat{\theta}_1 - \hat{\theta}_2 x_N \end{bmatrix}, \\ s^2 &= \frac{1}{N-2} \sum_{i=1}^N (y_i - \hat{\theta}_1 - \hat{\theta}_2 x_i)^2 = \\ &= \frac{1}{N-2} \sum_{i=1}^N \left(y_i - \frac{1}{N} \sum_{j=1}^N y_j - x_i \sum_{j=1}^N x_j y_j \left(\sum_{j=1}^N x_j^2 \right)^{-1} \right)^2 \quad (1.38) \end{aligned}$$

В качестве примера расчетов по формулам (1.25), (1.26) и (1.38) рассмотрим следующий случай. Предположим, что истинная функция регрессии имеет вид $\eta(x) = 1 + 2x$, $N = 6$, точки измерений имеют значения $x_1 = -1$, $x_2 = -0,6$, $x_3 = -0,2$, $x_4 = 0,2$, $x_5 = 0,6$, $x_6 = 1$. Значения ошибок измерений e_i выберем из таблицы нормально распределенных случайных величин с дисперсией $\sigma^2 = 1$ (табл. 4 из приложения 2). Пусть

$$\begin{aligned} e_1 &= -0,49, & e_2 &= 1,68, & e_3 &= -0,06, \\ e_4 &= -1,23, & e_5 &= -0,49, & e_6 &= 0,86. \end{aligned}$$

Таким образом, считаем, что результатами эксперимента являются числа

$$\begin{aligned} y_1 &= -1 - 0,49 = -1,49, & y_2 &= -0,2 + 1,68 = 1,48, \\ y_3 &= 0,6 - 0,06 = 0,54, \\ y_4 &= 1,4 - 1,23 = 0,17, & y_5 &= 2,2 - 0,49 = 1,71, \\ y_6 &= 3 + 0,86 = 3,86. \end{aligned}$$

Построим МНК-оценку $\hat{\theta} = (\hat{\theta}_1, \hat{\theta}_2)^T$ параметров регрессии $\theta_1 + \theta_2 x$. Имеем

$$F^T F = \begin{bmatrix} 1 & 1 & 1 & 1 & 1 & 1 \\ -1 & -0,6 & -0,2 & 0,2 & 0,6 & 1 \end{bmatrix} \begin{bmatrix} 1 & -1 \\ 1 & -0,6 \\ 1 & -0,2 \\ 1 & 0,2 \\ 1 & 0,6 \\ 1 & 1 \end{bmatrix} = \begin{bmatrix} 6 & 0 \\ 0 & 2,8 \end{bmatrix};$$

$$(F^T F)^{-1} = \begin{bmatrix} 0,17 & 0 \\ 0 & 0,36 \end{bmatrix};$$

$$F^T Y = \begin{bmatrix} 1 & 1 & 1 & 1 & 1 & 1 \\ -1 & -0,6 & -0,2 & 0,2 & 0,6 & 1 \end{bmatrix} \begin{bmatrix} -1,49 \\ 1,48 \\ 0,54 \\ 0,17 \\ 1,71 \\ 3,86 \end{bmatrix} = \begin{bmatrix} 6,27 \\ 5,41 \end{bmatrix};$$

$$\hat{\theta} = \begin{bmatrix} 0,17 & 0 \\ 0 & 0,36 \end{bmatrix} \begin{bmatrix} 6,27 \\ 5,41 \end{bmatrix} = \begin{bmatrix} 0,17 \times 6,27 \\ 0,36 \times 5,41 \end{bmatrix} = \begin{bmatrix} 1,07 \\ 1,95 \end{bmatrix}.$$

Вычислим теперь s^2 (несмещенную оценку $\sigma^2 = 1$):

$$(Y - F\hat{\theta}) =$$

$$= \begin{bmatrix} -1,49 \\ 1,48 \\ 0,54 \\ 0,17 \\ 1,71 \\ 3,86 \end{bmatrix} - \begin{bmatrix} 1 & -1 \\ 1 & -0,6 \\ 1 & -0,2 \\ 1 & 0,2 \\ 1 & 0,6 \\ 1 & 1 \end{bmatrix} \begin{bmatrix} 1,07 \\ 1,95 \end{bmatrix} = \begin{bmatrix} -1,49 - 1,07 + 1,95 \\ 1,48 - 1,07 + 0,6 \cdot 1,95 \\ 0,54 - 1,07 + 0,2 \cdot 1,95 \\ 0,17 - 1,07 - 0,2 \cdot 1,95 \\ 1,71 - 1,07 - 0,6 \cdot 1,95 \\ 3,86 - 1,07 - 1,95 \end{bmatrix} = \begin{bmatrix} -0,61 \\ 1,58 \\ -0,14 \\ -1,29 \\ -0,53 \\ 0,84 \end{bmatrix}.$$

$$\begin{aligned} (Y - F\hat{\theta})^T (Y - F\hat{\theta}) &= 0,37 + 2,50 + 0,02 + 1,66 + 0,28 + 0,70 = 5,53, \\ s^2 &= 5,53 / (6 - 2) = 5,53 / 4 = 1,38. \end{aligned}$$

1.5. Принцип максимального правдоподобия в классической линейной регрессионной модели. Для приведенных выше результатов не требуется задание типа распределения вектора ошибок измерений ε , а необходимо лишь задание среднего и дисперсионной матрицы этого вектора. В этом и следующем пунктах будет предполагаться, что случайный вектор ε имеет нормальное распределение $\varepsilon \sim N(0, \sigma^2 I_N)$, где обозначение $\varepsilon \sim N(a, \Sigma)$ есть краткая запись того, что случайный вектор ε имеет многомерное нормальное распределение с вектором средних a и дисперсионной матрицей Σ . Плотность этого распределения в случае невырожденности матрицы Σ равна

$$f(x; a, \Sigma) = (2\pi)^{-N/2} [\det \Sigma]^{-1/2} \exp \left\{ -\frac{1}{2} (x - a)^T \Sigma^{-1} (x - a) \right\};$$

в частном случае $a = 0, \Sigma = \sigma^2 I_N$

$$f(x; 0, \sigma^2 I_N) = (2\pi\sigma^2)^{-N/2} \exp \{ -x^T x / (2\sigma^2) \}.$$

При рассмотрении линейных регрессионных моделей вектор Y представляет собой выборку из распределения с некоторой плотностью $L(y, \theta)$, зависящей от $y \in \mathbb{R}^N$ и неизвестных параметров. Функцию $L(Y, \theta)$ как функцию от θ в математической статистике называют *функцией правдоподобия*, а значение $\theta = \hat{\theta}$, для которого функция правдоподобия принимает максимальное значение, — *оценкой максимального правдоподобия*.

Теорема 1.3. Пусть задана невырожденная классическая линейная регрессионная модель

$$Y = F\theta + \varepsilon, \quad \varepsilon \sim N(0, \sigma^2 I_N). \quad (1.39)$$

Тогда МНК-оценка (1.20) является также оценкой максимального правдоподобия параметров θ , а статистика $\hat{s}_N^2 = (N - m)s^2/N$ (величина s^2 определяется по формуле (1.35)) — оценкой максимального правдоподобия параметра σ^2 .

Доказательство. Поскольку из (1.39) следует, что $Y \sim N(F\theta, \sigma^2 I_N)$, функция максимального правдоподобия параметров θ, σ^2 имеет вид

$$L(Y, \theta, \sigma^2) = (2\pi\sigma^2)^{-N/2} \exp \left\{ -\frac{1}{2\sigma^2} (Y - F\theta)^T (Y - F\theta) \right\}, \quad (1.40)$$

поэтому можно решать задачу максимизации по θ, σ^2 функции

$$\ln L(Y, \theta, \sigma^2) = -\frac{N}{2} \ln(2\pi\sigma^2) - \frac{1}{2\sigma^2} (Y - F\theta)^T (Y - F\theta) \quad (1.41)$$

(логарифма правдоподобия). Дифференцируя (1.41) по θ и σ^2 (дифференцирование по θ производится так же, как и при доказательстве леммы 1.1) и приравнявая производные нулю, получаем

$$\sigma^2 \frac{\partial \ln L}{\partial \theta} = F^T (Y - F\theta) = 0, \quad (1.42)$$

$$2\sigma^2 \frac{\partial \ln L}{\partial (\sigma^2)} = -N + \frac{1}{\sigma^2} (Y - F\theta)^T (Y - F\theta) = 0. \quad (1.43)$$

Система уравнений (1.42) не что иное, как система нормальных уравнений (1.9), решение которой есть оценка МНК (1.20). Решением (1.43) является $\hat{s}_N$. Теорема доказана.

Отметим, что оценка максимального правдоподобия $\hat{s}_N^2$ для σ^2 является смещенной (т. е. $E\hat{s}_N^2 \neq \sigma^2$), но при $N \rightarrow \infty$ — асимптотически несмещенной (т. е. $E\hat{s}_N^2 \rightarrow \sigma^2$, $N \rightarrow \infty$).

Пример 1.8. Предположим, что функция регрессии — константа, т. е. результаты измерений представимы в виде $y_i = \theta + \varepsilon_i$ (функция регрессии зависит от одной переменной, которая тождественно равна единице). Пусть, как обычно, $E\varepsilon_i = 0$ при $i \neq j$ и $E\varepsilon_i^2 = \sigma^2$. В данном случае матрица F — это столбец, состоящий из единиц:

$$F^T F = N, \quad (F^T F)^{-1} = N^{-1}, \quad F^T Y = \sum_{i=1}^N y_i, \quad \hat{\theta} = \frac{1}{N} \sum_{i=1}^N y_i = \bar{y}.$$

Следовательно, МНК-оценка параметра θ — это выборочное среднее.

Несмещенной оценкой для дисперсии является

$$s^2 = \frac{1}{N-1} \sum_{i=1}^N (y_i - \bar{y})^2.$$

Из теоремы 1.3 вытекает классический результат о том, что при наличии одномерной повторной выборки из нормального распределения с неизвестными средним θ и дисперсией σ^2 оценкой максимального правдоподобия для θ является выборочное среднее $\bar{y}$, а для σ^2 — выборочная дисперсия $N^{-1} \sum_{i=1}^N (y_i - \bar{y})^2$.

1.6. Проверка гипотез о коэффициентах регрессии. В данном пункте показано, как для регрессии (1.39) проверять гипотезу $H_0: \theta_i = a$ о равенстве заданному числу некоторого θ_i при фиксированном i ($1 \leq i \leq m$) при альтернативе $H_1: \theta_i \neq a$. В случае $a = 0$ гипотеза H_0 называется *гипотезой значимости параметра* θ_i .

В рассматриваемой модели

$$\hat{\theta} = (\hat{\theta}_1, \dots, \hat{\theta}_m)^T \sim N(\theta, \sigma^2 (F^T F)^{-1}),$$

поэтому при справедливости гипотезы H_0 имеем

$$\hat{\theta}_i \sim N(a, D\hat{\theta}_i),$$

где $D\hat{\theta}_i = \sigma^2 [(F^T F)^{-1}]_{ii}$; A_{ii} есть i -й диагональный элемент матрицы A .

Если значение σ^2 известно, то при справедливости гипотезы H_0 статистика

$$u = (\hat{\theta}_i - a) / \sqrt{D\hat{\theta}_i} \quad (1.44)$$

имеет нормальное распределение $N(0, 1)$ и может быть выбрана в качестве тестовой статистики для проверки указанной гипотезы.

Если величина σ^2 неизвестна, то она может быть оценена статистикой s^2 . Оказывается, что в этом случае при справедливости гипотезы H_0 аналогом статистики (1.44) является статистика

$$t = (\theta_1 - a) / [s \sqrt{(F'F)^{-1}}_{11}], \quad (1.45)$$

которая имеет стандартное распределение, носящее название *t-распределения Стьюдента с $N - m$ степенями свободы*. Указанный факт вытекает из результатов п. 3.1.

Для проверки гипотезы H_0 следует использовать таблицы 1 и 2 приложения 2, в которых имеются значения квантилей нормального распределения и *t-распределения Стьюдента*.

Пусть зафиксировано некоторое число α ($0 < \alpha < 1$), определяющее уровень доверия приведенных процедур проверки гипотезы H_0 . Гипотеза H_0 должна быть отвергнута, если модуль величины u , определяемой по формуле (1.44), превышает модуль $\alpha/2$ -квантиля стандартного нормального распределения $N(0, 1)$ или (при использовании статистики (1.45)) если $|t|$ превышает модуль $\alpha/2$ -квантиля *t-распределения Стьюдента с $N - m$ степенями свободы*. Для приведенных процедур вероятность того, что гипотеза H_0 будет отвергнута при условии, что она верна, равна α .

Пример 1.9 (продолжение примера 1.7). Проверим гипотезу $H_0: \theta_1 = 0$ по результатам расчетов параметров регрессии $1 + 2x$, приведенных в примере 1.7. а) Пусть $\alpha = 0,05$ и $\sigma^2 = 1$. По формуле (1.44) рассчитаем статистику

$$u = 1,07 / \sqrt{0,17} = 2,60.$$

Значение модуля $\alpha/2$ -квантиля нормального распределения $N(0, 1)$ равно 1,96, т. е. меньше, чем $|u|$. Следовательно, на уровне значимости $1 - \alpha = 0,95$ при $\sigma^2 = 1$ гипотеза $H_0: \theta_1 = 0$ должна быть отвергнута. б) Пусть теперь σ^2 неизвестно. По формуле (1.45) рассчитаем статистику t :

$$t = 1,07 / (s \sqrt{0,17}) = 2,6 / \sqrt{1,38} = 2,21.$$

Значение модуля $\alpha/2$ -квантиля ($\alpha = 0,05$) *t-распределения с $N - m = 4$ степенями свободы* равно 2,78, и поэтому на уровне значимости $1 - \alpha = 0,95$ гипотеза $H_0: \theta_1 = 0$ должна быть принята (несмотря на то что истинное значение θ_1 равно 1).

1.7. Обобщенный метод наименьших квадратов. *Обобщенной линейной регрессионной моделью* называется модель $(Y, F\theta, \sigma^2 W)$, где W — известная положительно определенная $N \times N$ -матрица ($W > 0$), а параметр σ^2 может быть неизвестен.

Одной из типичных ситуаций, когда используется обобщенная линейная регрессионная модель, является такая ситуация, когда в классической модели измерения в точках повторяются. Действительно, пусть имеются результаты измерений

$$y_i = \theta' x_{i1} + \varepsilon_i, \quad (1.45)$$

где $x_{(j)} = (x_{j1}, \dots, x_{jm})^T$ — точки проведения измерений, e_j — некоррелированные случайные ошибки с одинаковой дисперсией σ^2 , и пусть среди точек $x_{(j)}$ ($j = 1, \dots, N$) лишь M ($M \leq N$) точек различные. Обозначим эти точки через $\bar{x}_{(1)}, \dots, \bar{x}_{(M)}$, а через r_l — число измерений в точке $\bar{x}_{(l)}$ ($r_l \geq 1, \sum_{l=1}^M r_l = N$).

Усредняя результаты различных измерений в точках $x_{(j)}$, получаем модель

$$\bar{y}_l = \theta^T \bar{x}_{(l)} + \bar{e}_l, \quad l = 1, \dots, M, \quad (1.47)$$

где $\bar{y}_l$ — средние арифметические результатов наблюдений y_j в точках $\bar{x}_{(l)}$; $\bar{e}_l$ — средние арифметические ошибок измерений e_j ; $E\bar{e}_l = 0$; $E\bar{e}_l \bar{e}_k = 0$ ($l \neq k$); $E\bar{e}_l^2 = \sigma^2/r_l$. Модель (1.47) является обобщенной линейной регрессионной моделью $(Y, F\theta, \sigma^2 W)$, где W — диагональная $M \times M$ -матрица с элементами $1/r_l$ на главной диагонали.

Пусть имеется обобщенная линейная регрессионная модель $(Y, F\theta, \sigma^2 W)$. Поскольку $W > 0$, то в силу теоремы 1.3 из приложения 1 существует такая невырожденная $N \times N$ -матрица V , что $W = VV^T$. Полагая $\bar{Y} = V^{-1}Y$, $\bar{F} = V^{-1}F$, $\bar{e} = V^{-1}e$ и умножая обе части (1.11) слева на V^{-1} , получаем, что модель $(Y, F\theta, \sigma^2 W)$ эквивалентна модели $(\bar{Y}, \bar{F}\theta, \sigma^2 I_N)$, так как

$$\begin{aligned} D\bar{e} &= E\bar{e}\bar{e}^T = EV^{-1}e e^T (V^{-1})^T = V^{-1}\sigma^2 W (V^{-1})^T = \\ &= \sigma^2 V^{-1} V V^T (V^T)^{-1} = \sigma^2 I_N. \end{aligned}$$

МНК-оценка для модели $(\bar{Y}, \bar{F}\theta, \sigma^2 I_N)$ записывается в виде

$$\begin{aligned} \bar{\theta} &= (\bar{F}^T \bar{F})^{-1} \bar{F}^T \bar{Y} = (F^T (V V^T)^{-1} F)^{-1} F^T (V V^T)^{-1} Y = \\ &= (F^T W^{-1} F)^{-1} F^T W^{-1} Y. \end{aligned}$$

По определению оценка

$$\hat{\theta} = (F^T W^{-1} F)^{-1} F^T W^{-1} Y \quad (1.48)$$

называется *обобщенной МНК-оценкой*. В силу леммы 1.1 эта оценка минимизирует выражение

$$(\bar{Y} - \bar{F}\hat{\theta})^T (\bar{Y} - \bar{F}\hat{\theta}) = (Y - F\hat{\theta})^T W^{-1} (Y - F\hat{\theta}). \quad (1.49)$$

В частном случае диагональной матрицы W выражение (1.49) называется *взвешенной суммой квадратов отклонений*.

Дисперсионная матрица обобщенной МНК-оценки (1.48) равна, как нетрудно видеть,

$$D\hat{\theta} = \sigma^2 (F^T W^{-1} F)^{-1}. \quad (1.50)$$

Действительно, используя лемму 1.3, имеем

$$\begin{aligned} D\hat{\theta} &= (F^T W^{-1} F)^{-1} F^T W^{-1} D Y [(F^T W^{-1} F)^{-1} F^T W^{-1}]^T = \\ &= \sigma^2 (F^T W^{-1} F)^{-1} F^T W^{-1} W W^{-1} F (F^T W^{-1} F)^{-1} = \sigma^2 (F^T W^{-1} F)^{-1}. \end{aligned}$$

В силу теоремы Гаусса — Маркова (теорема 1.1) обобщенная МНК-оценка (1.48) является НЛН-оценкой для модели $(Y, F\theta, \sigma^2/W)$ и, следовательно, для модели $(Y, F\theta, \sigma^2 W)$.

Пример 1.10 (продолжение примера 1.7). Предположим, что истинная функция регрессии $\eta(x) = 1 + 2x$, и шесть измерений проводится в трех точках: $x_1 = -1$ (трижды), $x_2 = 0$ и $x_3 = 1$ (дважды). Будем считать, что дисперсия ошибок измерений σ^2 равна единице, а сами эти ошибки те же, что и в примере 1.7: $-0,49, 1,68, -0,06$ (при измерениях в точке x_1); $-0,49, 0,86$ (при измерениях в точке x_3); $-1,23$ (при измерении в точке x_2).

Усредняя результаты измерений в точках x_1 и x_3 , получаем

$$y_1 = -1 + \frac{1}{3}(-0,49 + 1,86 - 0,06) = -0,56,$$

$$y_2 = 1 - 1,23 = -0,23,$$

$$y_3 = 3 + \frac{1}{2}(-0,49 + 0,86) = 3,19.$$

Дисперсии ошибок при измерении y_1, y_2, y_3 равны соответственно $\sigma^2/3, \sigma^2, \sigma^2/2$. Таким образом, в рассматриваемой модели $(Y, F\theta, \sigma^2 W)$ имеем

$$Y = \begin{bmatrix} -0,56 \\ -0,23 \\ 3,19 \end{bmatrix}, \quad F = \begin{bmatrix} 1 & -1 \\ 1 & 0 \\ 1 & 1 \end{bmatrix}, \quad W = \begin{bmatrix} 1/3 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1/2 \end{bmatrix}.$$

Вычислим обобщенную МНК-оценку (1.48):

$$\begin{aligned} W^{-1} &= \begin{bmatrix} 3 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 2 \end{bmatrix}, \\ F^T W^{-1} F &= \begin{bmatrix} 1 & 1 & 1 \\ -1 & 0 & 1 \end{bmatrix} \begin{bmatrix} 3 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 2 \end{bmatrix} \begin{bmatrix} 1 & -1 \\ 1 & 0 \\ 1 & 1 \end{bmatrix} = \\ &= \begin{bmatrix} 3 & 1 & 2 \\ -3 & 0 & 2 \end{bmatrix} \begin{bmatrix} 1 & -1 \\ 1 & 0 \\ 1 & 1 \end{bmatrix} = \begin{bmatrix} 5 & -1 \\ -1 & 5 \end{bmatrix}, \\ (F^T W^{-1} F)^{-1} &= \begin{bmatrix} 5/29 & 1/29 \\ 1/29 & 6/29 \end{bmatrix}, \\ F^T W^{-1} Y &= \begin{bmatrix} 3 & 1 & 2 \\ -3 & 0 & 2 \end{bmatrix} \begin{bmatrix} -0,56 \\ -0,23 \\ 3,19 \end{bmatrix} = \begin{bmatrix} 4,47 \\ 8,03 \end{bmatrix}, \\ \hat{\theta} &= \begin{bmatrix} 0,17 & 0,034 \\ 0,034 & 0,207 \end{bmatrix} \begin{bmatrix} 4,47 \\ 8,03 \end{bmatrix} = \begin{bmatrix} 0,77 + 0,28 \\ 0,16 + 1,66 \end{bmatrix} = \begin{bmatrix} 1,05 \\ 1,82 \end{bmatrix}. \end{aligned}$$

В предположении $\sigma^2 = 1$ проверим, согласно процедуре п. 1.6, гипотезу $H_0: \theta_2 = 2$ (т. е. проверим, значительно ли отклонение

оценки $\hat{\theta}_2 = 1,82$ от истинного значения $\theta_2 = 2$). Дисперсия оценки $\hat{\theta}_2$ равна

$$D\hat{\theta}_2 = \sigma^2 [(F^T W^{-1} F)^{-1}]_{22} = \frac{6}{29} = 0,207.$$

Поэтому величина (1.44) для $a = 2$ равна

$$u = \frac{1,82 - 2}{\sqrt{0,207}} = -0,40, \quad |u| \approx 0,4,$$

и, следовательно, для любого разумного значения уровня доверия $1 - \alpha$ гипотеза H_0 должна быть принята. Действительно, пусть $\alpha = 0,05$, тогда модуль $\alpha/2$ -квантиля нормального распределения $N(0, 1)$ равен 1,96, т. е. больше, чем $|u|$.

1.8. Последствия неправильного задания дисперсионной матрицы ошибок измерений. Предположим, что для оценки параметров обобщенной линейной регрессионной модели $(Y, F\theta, \sigma^2 W)$ используется оценка

$$\hat{\theta}(A) = (F^T A^{-1} F)^{-1} F^T A^{-1} Y,$$

где $A > 0$ — некоторая матрица, обычно являющаяся некоторым приближением к W .

Как нетрудно проверить, эта оценка является несмещенной и имеет дисперсионную матрицу

$$D[\hat{\theta}(A)] = \sigma^2 (F^T A^{-1} F)^{-1} F^T A^{-1} W A^{-1} F (F^T A^{-1} F)^{-1}.$$

Поскольку $\hat{\theta}(W)$ является НЛН-оценкой,

$$D[\hat{\theta}(A)] \geq D[\hat{\theta}(W)].$$

В наиболее интересном частном случае $A = I_N$, т. е. когда используется оценка (1.20), потеря эффективности определяется формулой

$$D[\hat{\theta}(I_N)] - D[\hat{\theta}(W)] = \sigma^2 (UF^T - SF^T W^{-1}) W (FU - W^{-1} FS) = \\ = \sigma^2 (UF^T W F U - S) \geq 0,$$

где $S = (F^T W^{-1} F)^{-1}$, $U = (F^T F)^{-1}$.

Отсюда, в частности, следует, что точность оценки (1.20) такая же, как и оценки (1.48), в том и только том случае, когда выполнено

$$UF^T = SF^T W^{-1}.$$

Рассмотрим теперь последствия задания матрицы $A = I_N$ вместо W при оценивании σ^2 , т. е. при использовании формулы (1.35) для оценивания σ^2 в модели $(Y, F\theta, \sigma^2 W)$.

Имеем $s^2 = e^T e / (N - m)$, где

$$\bar{e} = Y - FUF^T Y = GY, \quad G = I_N - FUF^T, \quad G^2 = G, \quad \text{tr } G = N - m$$

(см. п. 1.4). Аналогично доказательству теоремы 1.2 получаем

$$\begin{aligned} \mathbf{E} \mathbf{e}^T \mathbf{e} &= \mathbf{E} \mathbf{Y}^T \mathbf{G}^2 \mathbf{Y} = \mathbf{E} \mathbf{Y}^T \mathbf{G} \mathbf{Y} = \mathbf{E} (\mathbf{F} \boldsymbol{\theta} + \mathbf{e})^T \mathbf{G} (\mathbf{F} \boldsymbol{\theta} + \mathbf{e}) = \\ &= \boldsymbol{\theta}^T \mathbf{F}^T \mathbf{G} \mathbf{F} \boldsymbol{\theta} + \mathbf{E} \mathbf{e}^T \mathbf{G} \mathbf{e} = \boldsymbol{\theta}^T \mathbf{F}^T \mathbf{F} \boldsymbol{\theta} - \boldsymbol{\theta}^T \mathbf{F}^T \mathbf{F} \boldsymbol{\theta} + \mathbf{E} \operatorname{tr} \mathbf{G} \mathbf{e} \mathbf{e}^T = \\ &= \operatorname{tr} \mathbf{G} \mathbf{E} \mathbf{e} \mathbf{e}^T = \sigma^2 \operatorname{tr} \mathbf{G} \mathbf{W} = \sigma^2 \operatorname{tr} \mathbf{W} - \sigma^2 \operatorname{tr} \mathbf{F} \mathbf{U} \mathbf{F}^T \mathbf{W} = \\ &= \sigma^2 [\operatorname{tr} \mathbf{W} - \operatorname{tr} (\mathbf{F}^T \mathbf{W} \mathbf{F} \mathbf{U})]. \end{aligned}$$

Следовательно,

$$\begin{aligned} \mathbf{E} s^2 &= \frac{\mathbf{E} \mathbf{e}^T \mathbf{e}}{N-m} = \frac{\sigma^2 [\operatorname{tr} \mathbf{G} + \operatorname{tr} \mathbf{G} (\mathbf{W} - \mathbf{I}_N)]}{N-m} = \\ &= \sigma^2 + \frac{\sigma^2}{N-m} \operatorname{tr} \mathbf{G} (\mathbf{W} - \mathbf{I}_N) = \sigma^2 - \frac{\sigma^2}{N-m} [\operatorname{tr} (\mathbf{F}^T \mathbf{W} \mathbf{F} \mathbf{U}) - m]. \end{aligned}$$

Таким образом, s^2 не является, вообще говоря, несмещенной оценкой для σ^2 .

1.9. Вычисление МНК-оценок для моделей неполного ранга. Говорят, что линейная регрессионная модель $(Y, F\boldsymbol{\theta}, \Sigma)$ имеет неполный ранг, если $\operatorname{rang} F = p < m$. Линейные регрессионные модели неполного ранга распространены, в частности, в задачах дисперсионного анализа. Для указанных моделей система нормальных уравнений имеет неединственное решение. Опишем способ построения МНК-оценок, т. е. решения системы нормальных уравнений

$$\mathbf{F}^T \mathbf{F} \boldsymbol{\theta} = \mathbf{F}^T \mathbf{Y}. \quad (1.51)$$

Предположим, что ранг матрицы $\mathbf{F}^T \mathbf{F}$ равен p ($\operatorname{rang} \mathbf{F}^T \mathbf{F} = p$), где $0 \leq p \leq m$. Не ограничивая общности, можно считать, что первые p столбцов $X_1, \dots, X_p$ $N \times m$ -матрицы $\mathbf{F}$ линейно независимы. Имеем $\mathbf{F} = (\mathbf{F}_{(1)}, \mathbf{F}_{(2)})$, $\boldsymbol{\theta} = (\boldsymbol{\theta}_{(1)}, \boldsymbol{\theta}_{(2)})^T$, где

$$\mathbf{F}_{(1)} = (X_1, \dots, X_p), \quad \mathbf{F}_{(2)} = (X_{p+1}, \dots, X_m)$$

суть матрицы размера $N \times p$ и $N \times (m-p)$ соответственно,

$$\boldsymbol{\theta}_{(1)} = (\theta_1, \dots, \theta_p)^T, \quad \boldsymbol{\theta}_{(2)} = (\theta_{p+1}, \dots, \theta_m)^T.$$

Так как $\operatorname{rang} F = \operatorname{rang} F_{(1)} = p$, то столбцы матрицы $\mathbf{F}_{(2)}$ линейно зависят от столбцов $\mathbf{F}_{(1)}$, т. е. существует такая $p \times (m-p)$ -матрица $\mathbf{L}$, что $\mathbf{F}_{(2)} = \mathbf{F}_{(1)} \mathbf{L}$. Система нормальных уравнений (1.51) записывается в виде

$$\begin{bmatrix} \mathbf{F}_{(1)}^T \mathbf{F}_{(1)} & \mathbf{F}_{(1)}^T \mathbf{F}_{(1)} \mathbf{L} \\ \mathbf{L}^T \mathbf{F}_{(1)}^T \mathbf{F}_{(1)} & \mathbf{L}^T \mathbf{F}_{(1)}^T \mathbf{F}_{(1)} \mathbf{L} \end{bmatrix} \cdot \begin{bmatrix} \boldsymbol{\theta}_{(1)} \\ \boldsymbol{\theta}_{(2)} \end{bmatrix} = \begin{bmatrix} \mathbf{F}_{(1)}^T \mathbf{Y} \\ \mathbf{L}^T \mathbf{F}_{(1)}^T \mathbf{Y} \end{bmatrix}. \quad (1.52)$$

Первые p строк этой системы уравнений имеют вид

$$\mathbf{F}_{(1)}^T \mathbf{F}_{(1)} \boldsymbol{\theta}_{(1)} + \mathbf{F}_{(1)}^T \mathbf{F}_{(1)} \mathbf{L} \boldsymbol{\theta}_{(2)} = \mathbf{F}_{(1)}^T \mathbf{Y}. \quad (1.53)$$

Остальные $m-p$ уравнений (1.52) — линейная комбинация уравнений (1.53); поэтому любое решение системы уравнений (1.53) представляет собой решение системы (1.51).

Поскольку $\text{rang } F_{(1)} = p$, у матрицы $F_{(1)}^T F_{(1)}$ существует обратная. Домножим обе части (1.53) слева на $(F_{(1)}^T F_{(1)})^{-1}$:

$$\theta_{(1)} = (F_{(1)}^T F_{(1)})^{-1} F_{(1)}^T Y - L\theta_{(2)}. \quad (1.54)$$

Поскольку $(m - p)$ -мерный вектор $\theta_{(2)}$ может быть выбран произвольным образом, то формула (1.54) определяет $(m - p)$ -мерное многообразие решений системы нормальных уравнений (1.51).

Пример 1.11. Пусть функция регрессии имеет вид (1.17), т. е. $\eta(x) = \theta_1 x_1 + \theta_2 x_2 + \theta_3 x_3$, а измерения проводятся в трех точках: $(1, -1, 0)^T$, $(0, 1, -1)^T$, $(-1, 0, 1)^T$. В интерпретации схемы взвешивания трех предметов A, B, C на двухчашечных весах (см. пример 1.4) это соответствует тому, что при первом взвешивании на левой чашке находился предмет A , на правой — предмет B ; при втором взвешивании на левой — предмет B , на правой — предмет C ; при третьем — соответственно C и A . Матрица F имеет вид

$$F = \begin{bmatrix} 1 & -1 & 0 \\ 0 & 1 & -1 \\ -1 & 0 & 1 \end{bmatrix}.$$

Имеем

$$F^T F = \begin{bmatrix} 2 & -1 & -1 \\ -1 & 2 & -1 \\ -1 & -1 & 2 \end{bmatrix}, \quad FY = \begin{bmatrix} y_1 - y_2 \\ y_2 - y_1 \\ y_3 - y_2 \end{bmatrix};$$

поэтому система нормальных уравнений имеет вид

$$\begin{aligned} 2\theta_1 - \theta_2 - \theta_3 &= y_1 - y_2, \\ -\theta_1 + 2\theta_2 - \theta_3 &= y_2 - y_1, \\ -\theta_1 - \theta_2 + 2\theta_3 &= y_3 - y_2. \end{aligned}$$

Третье уравнение является следствием первых двух и получается после их сложения и умножения обеих частей полученного уравнения на -1 . Следовательно, значение θ_3 можно выбрать произвольным, а первые два уравнения преобразовать так, чтобы θ_1 и θ_2 выразить через y_1, y_2, y_3 и θ_3 .

Из второго уравнения получаем

$$\theta_1 = 2\theta_2 - \theta_3 - y_2 + y_1$$

и подставляем последнее в первое уравнение системы:

$$4\theta_2 - 2\theta_3 - 2y_2 + 2y_1 - \theta_2 - \theta_3 = y_1 - y_2.$$

Окончательно имеем

$$\theta_2 = \frac{1}{3}(-y_1 + 2y_2 - y_3) + \theta_3,$$

$$\begin{aligned} \theta_1 &= \frac{1}{3}(-2y_1 + 4y_2 - 2y_3 + 4\theta_3 - 3\theta_3 - 3y_2 + 3y_1) - \\ &= \frac{1}{3}(y_1 + y_2 - 2y_3) + \theta_3. \end{aligned}$$

Построим теперь МНК-оценку, используя (1.54). Положим

$$\theta_{(1)} = \begin{bmatrix} \theta_1 \\ \theta_2 \end{bmatrix}, \quad F_{(1)} = \begin{bmatrix} 1 & -1 \\ 0 & 1 \\ -1 & 0 \end{bmatrix}, \quad F_{(2)} = \begin{bmatrix} 0 \\ -1 \\ 1 \end{bmatrix} = F_{(1)}L,$$

где $L = \begin{bmatrix} -1 \\ 1 \end{bmatrix}$. Тогда

$$F_{(1)}^T F_{(1)} = \begin{bmatrix} 1 & 0 & -1 \\ -1 & 1 & 0 \end{bmatrix} \begin{bmatrix} 1 & -1 \\ 0 & 1 \\ -1 & 0 \end{bmatrix} = \begin{bmatrix} 2 & -1 \\ -1 & 2 \end{bmatrix},$$

$$(F_{(1)}^T F_{(1)})^{-1} = \begin{bmatrix} 2/3 & 1/3 \\ 1/3 & 2/3 \end{bmatrix}, \quad F_{(1)}^T Y = \begin{bmatrix} y_1 - y_3 \\ y_2 - y_1 \end{bmatrix},$$

и по формуле (1.54) получаем

$$\theta_{(1)} = \begin{bmatrix} 2/3 & 1/3 \\ 1/3 & 2/3 \end{bmatrix} \begin{bmatrix} y_1 - y_3 \\ y_2 - y_1 \end{bmatrix} + \begin{bmatrix} \theta_3 \\ \theta_3 \end{bmatrix} = \begin{bmatrix} (1/3)(y_1 - 2y_3 + y_2) + \theta_3 \\ (1/3)(-y_1 - y_3 + 2y_2) + \theta_3 \end{bmatrix}.$$

Оба способа решения привели к одинаковому результату.

Если $\hat{\theta}$ — оценка параметров θ , то оценкой вектора значений функции регрессии $F\hat{\theta}$ является, конечно, $F\hat{\theta}$. Покажем, что при оценивании по МНК оценка вектора значений функции регрессии не зависит от выбора решения системы нормальных уравнений.

Лемма 1.5. Пусть $\text{rang } F = p \leq m$. Любые два решения $\hat{\theta}^{(1)}$ и $\hat{\theta}^{(2)}$ системы нормальных уравнений (1.51) приводят к одной и той же оценке вектора значений функции регрессии, т. е.

$$F\hat{\theta}^{(1)} = F\hat{\theta}^{(2)}$$

и, кроме того, $Q(\hat{\theta}^{(1)}) = Q(\hat{\theta}^{(2)})$, где

$$Q(\theta) = (Y - F\theta)^T (Y - F\theta). \quad (1.55)$$

Доказательство. Любое решение $\hat{\theta}$ системы уравнений (1.51) представимо в виде $\hat{\theta} = (\hat{\theta}_{(1)}, \hat{\theta}_{(2)})^T$, где $\hat{\theta}_{(2)}$ произвольно, а $\hat{\theta}_{(1)}$ определяется по формуле (1.54). Отсюда получаем

$$\begin{aligned} F\hat{\theta} &= (F_{(1)}, F_{(1)}L) \begin{bmatrix} (F_{(1)}^T F_{(1)})^{-1} F_{(1)}^T Y - L\theta_{(2)} \\ \theta_{(2)} \end{bmatrix} = \\ &= F_{(1)} (F_{(1)}^T F_{(1)})^{-1} F_{(1)}^T Y - F_{(1)}L\theta_{(2)} + F_{(1)}L\theta_{(2)}, \end{aligned}$$

т. е. $F\hat{\theta}$ однозначно определяется по матрице данных (1.1). Теперь очевидно, что так же определяется и функция $Q(\hat{\theta})$. Лемма доказана.

Из леммы 1.5 следует, в частности, что любое решение системы нормальных уравнений (1.51) доставляет минимум функции Q , т. е. является оценкой МНК. Обратное утверждение (любая оценка МНК является решением (1.51)) следует из леммы 1.1.

Другой способ представления решений системы уравнений (1.51) основан на использовании псевдообратных матриц (см. п. 1.6 приложения 1): любое решение системы нормальных

уравнений (1.51) может быть записано в виде

$$\hat{\theta} = (F^T F)^{-1} F^T Y, \quad (1.56)$$

где A^{-} — обобщенная обратная матрица для квадратной матрицы A , т. е. такая матрица, что $AA^{-}A = A$. Отметим, что (1.54) и (1.56) различаются только формой записи.

Для случая, когда $\text{rang } F = p < m$ несмещенно оценить все неизвестные параметры модели $(Y, F\theta, \Sigma)$ невозможно, но можно оценить некоторые линейные комбинации неизвестных параметров θ (параметрические функции). Любая параметрическая функция представима в виде $\tau = t^T \theta$, где t есть m -вектор.

Говорят, что параметрическая функция $\tau = t^T \theta$ оцениваема, если для нее существует линейная несмещенная оценка вида $\hat{\tau} = b^T Y$ ($b \in \mathbb{R}^n$).

Лемма 1.6. Для модели $(Y, F\theta, \Sigma)$ параметрическая функция $\tau = t^T \theta$ оцениваема тогда и только тогда, когда $t \in \mathcal{L}[F^T]$, т. е. t является линейной комбинацией столбцов матрицы F^T .

Доказательство. В силу несмещенности оценки $\hat{\tau} = b^T Y$ для всех θ имеем

$$t^T \theta = E \hat{\tau} = E b^T Y = b^T F \theta.$$

Это равносильно тому, что $t^T = b^T F$ или $t = F^T b$ (см. доказательство леммы 1.2). Лемма доказана.

Теорема 1.4. Пусть задана модель $(Y, F\theta, \sigma^2 I_N)$, $\text{rang } F = p \leq m$ и $\hat{\theta}$ — произвольное решение системы нормальных уравнений (1.51). Тогда, если параметрическая функция $\tau = t^T \theta$ оцениваема, то

а) вид оценки $t^T \hat{\theta}$ не зависит от выбора решения системы нормальных уравнений;

б) $t^T \hat{\theta}$ является НЛН-оценкой для τ .

Доказательство. Пусть параметрическая функция $\tau = t^T \theta$ оцениваема. Из леммы 1.6 вытекает, что существует такой N -вектор b , для которого $t^T = b^T F$, откуда $t^T \hat{\theta} = b^T F \hat{\theta}$. В силу леммы 1.5 правая часть последнего равенства определена однозначно. Утверждение второй части теоремы доказывается точно так же, как теорема 1.1. Теорема доказана.

Приведенная теорема легко обобщается на тот случай, когда необходимо оценить векторную параметрическую функцию $\tau = T\theta$, где T — такая $q \times m$ -матрица, что ее строки являются линейными комбинациями строк матрицы F (в этом случае, согласно лемме 1.6, существует линейная несмещенная оценка для τ , т. е. τ оцениваема).

Теорема 1.5 (обобщение теоремы 1.4). Пусть задана модель $(Y, F\theta, \sigma^2 I_N)$, $\text{rang } F = p \leq m$ и $\hat{\theta}$ — произвольное решение системы нормальных уравнений (1.51). Тогда, если векторная параметрическая функция $\tau = T\theta$ оцениваема, то оценка $\hat{\tau} = T\hat{\theta}$ определяется единственным образом и является НЛН-оценкой для τ .

Доказательство второй части теоремы аналогично доказательству теоремы 1.1, а первая часть (единственность τ) доказана выше.

1.10. Смещенное линейное оценивание. Оценка МНК $\hat{\theta}$ является наилучшей в классе линейных несмещенных оценок. Однако если отказаться от условия несмещенности, то можно получить оценки, которые в некоторых случаях обладают рядом преимуществ по сравнению с МНК-оценкой.

Рассмотрим невырожденную модель $(Y, F\theta, \sigma^2 I_N)$. Преобразуем функцию (1.55). Учитывая, что $F^T Y = F^T F \hat{\theta}$, где $\hat{\theta}$ есть МНК-оценка (1.20), получаем

$$Q(\theta) = (Y - F\theta)^T (Y - F\theta) = [Y - F\hat{\theta} - F(\theta - \hat{\theta})]^T \times \\ \times [Y - F\hat{\theta} - F(\theta - \hat{\theta})] = (Y - F\hat{\theta})^T (Y - F\hat{\theta}) - (Y - F\hat{\theta})^T F(\theta - \hat{\theta}) - \\ - (\theta - \hat{\theta})^T F^T (Y - F\hat{\theta}) + (\hat{\theta} - \theta)^T F^T F (\hat{\theta} - \theta) = Q_{\min} + a(\theta),$$

где

$$Q_{\min} = Q(\hat{\theta}) = (Y - F\hat{\theta})^T (Y - F\hat{\theta}), \quad a(\theta) = (\hat{\theta} - \theta)^T F^T F (\hat{\theta} - \theta).$$

Существует многообразие векторов θ , удовлетворяющих соотношению

$$Q(\theta) = Q_{\min} + a_0, \quad (1.57)$$

где $a_0 > 0$ — заданное число. При заданном соотношении (1.57) минимизация одного из возможных критериев качества оценки — квадрата ее длины (квадрат длины вектора θ по определению равен $\theta^T \theta$) — в силу известного в теории экстремальных задач метода множителей Лагранжа [2] ведет к задаче поиска минимума по θ функции

$$R(\theta) = \theta^T \theta + (1/k) [(\hat{\theta} - \theta)^T F^T F (\hat{\theta} - \theta) - a_0], \quad (1.58)$$

где $1/k$ — множитель Лагранжа. Функция $R(\theta)$ — квадратичная форма, и поэтому (см. доказательство леммы 1.1) для нахождения точки минимума этой функции достаточно решить систему уравнений $\nabla R(\theta) = 0$. Обозначим элементы матрицы $F^T F$ через m_{ij} ($i, j = 1, \dots, m$) и заметим, что $m_{ij} = m_{ji}$. Имеем

$$\begin{aligned} \frac{\partial R(\theta)}{\partial \theta_i} &= \frac{\partial}{\partial \theta_i} \sum_{l=1}^m \theta_l^2 + \frac{1}{k} \frac{\partial}{\partial \theta_i} \sum_{l,j=1}^m (\theta_l - \hat{\theta}_l) m_{lj} (\theta_l - \hat{\theta}_l) = \\ &= 2\theta_i + \frac{1}{k} \frac{\partial}{\partial \theta_i} \sum_{l=1}^m (\theta_l - \hat{\theta}_l)^2 m_{ll} + \frac{1}{k} \frac{\partial}{\partial \theta_i} \sum_{l \neq j}^m (\theta_l - \hat{\theta}_l) m_{lj} (\theta_l - \hat{\theta}_l) = \\ &= 2\theta_i + \frac{1}{k} 2m_{ii} (\theta_i - \hat{\theta}_i) + \frac{2}{k} \sum_{\substack{l=1 \\ l \neq i}}^m (\theta_l - \hat{\theta}_l) m_{li} = \\ &= 2 \left[\theta_i + \frac{1}{k} \sum_{l=1}^m (\theta_l - \hat{\theta}_l) m_{li} \right]. \end{aligned}$$

Следовательно,

$$\nabla R(\theta) = \left(\frac{\partial R(\theta)}{\partial \theta_1}, \dots, \frac{\partial R(\theta)}{\partial \theta_m} \right)^T = 2 \left[\theta + \frac{1}{k} F^T F (\theta - \hat{\theta}) \right].$$

Поскольку из вида МНК-оценки имеем

$$F^T F \hat{\theta} = (F^T F) (F^T F)^{-1} F^T Y = F^T Y,$$

система уравнений $\nabla R(\theta) = 0$ записывается в виде

$$\theta + \frac{1}{k} F^T F \theta = \frac{1}{k} F^T Y.$$

Преобразуя эту систему уравнений, получаем

$$(kI_m + F^T F) \theta = F^T Y,$$

$$\theta = (kI_m + F^T F)^{-1} F^T Y.$$

Полученная оценка

$$\hat{\theta}(k) = (F^T F + kI_m)^{-1} F^T Y \quad (1.59)$$

называется *гребневой*. Число k однозначно связано с числом a_0 соотношением (1.57):

$$Q(\hat{\theta}(k)) = Q_{\min} + a_0. \quad (1.60)$$

Легко понять, что при малых k оценки $\hat{\theta}(k)$ и $\hat{\theta}$ почти совпадают ($\hat{\theta}(k) \rightarrow \hat{\theta}$ с вероятностью 1 при $k \rightarrow 0$). Преимущества гребневой оценки (1.59) по сравнению со стандартной МНК-оценкой (1.20) проявляются в тех случаях, когда матрица $F^T F$ плохо обусловлена. Отметим, что в отличие от стандартной МНК-оценки гребневая оценка однозначно определяется вне зависимости от того, вырождена модель или нет.

Гребневая оценка (1.59) встречается в § 2 как оптимальная при наличии разного рода априорной информации о параметрах.

Другой распространенный тип смещенных линейных оценок — *сжимающие оценки* (называемые также *оценками Джеймса — Стейна*) $\hat{\theta}(\alpha) = \alpha \hat{\theta}$, где $\hat{\theta}$ — стандартная МНК-оценка (1.20), $0 < \alpha < 1$. Для невырожденной классической линейной регрессионной модели $(Y, F\theta, \sigma^2 I_N)$ для всех θ имеем

$$E\hat{\theta}(\alpha) = \alpha E\hat{\theta} = \alpha \theta,$$

$$D\hat{\theta}(\alpha) = \alpha^2 D\hat{\theta}.$$

Таким образом, оценка $\hat{\theta}(\alpha)$ смещенная (в среднем недооценивает вектор θ), но дисперсионная матрица оценки $\hat{\theta}(\alpha)$ всегда меньше, чем дисперсионная матрица МНК-оценки $\hat{\theta}$. Следовательно, разумный выбор близкого к единице числа α позволяет за счет некоторого смещения уменьшить дисперсии и ковариации оценок.

Упражнения.

1. В обозначениях примера 1.5 покажите, что оценки (1.22) и (1.23) некоррелированы (т. е. $E\theta_1\theta_2 = E\theta_1 E\theta_2$) тогда и только тогда, когда выполнено условие (1.24).

2. Пусть

$$y_j = \theta_1 + \theta_2 x_j + \theta_3 (3x_j - 2) + e_j, \quad j = 1, 2, 3, \\ e \sim (0, \sigma^2 I_3), \quad x_1 = -1, \quad x_2 = 0, \quad x_3 = 1.$$

Найдите МНК-оценки параметров θ_1 , θ_2 и покажите, что эти оценки совпадают с МНК-оценками, полученными при условии $\theta_3 = 0$.

3. Выберите из табл. 4 приложения 2 другие значения ошибок измерений и проведите вычисления, аналогичные проведенным в примере 1.7.

4. Пусть $Y \sim N(F\theta, \sigma^2 I_N)$, где $N \times m$ -матрица плапа F имеет ранг m , и пусть s^2 вычисляется по формуле (1.35), а матрица A — по формуле

$$A = \frac{1}{N - m + 2} (I_N - F(F^T F)^{-1} F^T).$$

Вычислите Ds^2 и $E[(Y^T A Y - \sigma^2)^2]$ и покажите, что

$$Ds^2 > E[(Y^T A Y - \sigma^2)^2],$$

т. е. что оценка $Y^T A Y$ для σ^2 имеет меньшую среднеквадратичную ошибку, чем s^2 .

5. Пусть $y_1 = \theta + e_1$, $y_2 = -\theta + e_2$,

$$e_1 \sim N(0, \sigma^2), \quad e_2 \sim N(0, 10\sigma^2), \quad E e_1 e_2 = 0$$

Запишите эту схему линейной регрессии в стандартном виде. Найдите обобщенную МНК-оценку для θ и дисперсию этой оценки.

6. Пусть

$$y_j \sim (\theta, w_j \sigma^2), \quad j = 1, \dots, N,$$

причем случайные величины $y_1, \dots, y_N$ некоррелированы. Найдите НЛН-оценку для θ .

7. Пусть

$$y_j \sim N(j\theta, j^2 \sigma^2), \quad j = 1, \dots, N,$$

причем случайные величины $y_1, \dots, y_N$ некоррелированы. Найдите НЛН-оценку для θ и покажите, что ее дисперсия равна σ^2/N .

8. Приведите пример, когда $W \neq I_N$, но дисперсионная матрица стандартной МНК-оценки совпадает с дисперсионной матрицей обобщенной МНК-оценки.

9. Показать, что для схемы регрессии $(Y, F\theta, \sigma^2 I_N)$ линейная форма $a^T \theta$ оцениваема тогда и только тогда, когда

$$a^T (F^T F)^- F^T F = a^T.$$

10. Покажите, что если в схеме регрессии $(Y, F\theta, \sigma^2 I_N)$ линейная форма $a^T \theta$ оцениваема, то

$$D[a^T \theta] = \sigma^2 a^T (F^T F)^- a.$$

11. Пусть имеется схема линейной регрессии $(Y, F\theta, \sigma^2 I_N)$. Выведите выражение для дисперсионной матрицы оцениваемой векторной параметрической функции

12. Для схемы регрессии, приведенной в примере 1.10, напишите выражение для гребисвой оценки. Используя (1.60), выведите уравнение связи между параметрами k и a_0 .

§ 2. Линейный регрессионный анализ при наличии априорной информации о параметрах

2.1. Основные виды априорной информации о параметрах. Обобщенная МНК-оценка (1.48) является НЛН-оценкой для модели $(Y, F\theta, \sigma^2 W)$. При этом источником информации о неизвестных параметрах являются результаты измерений Y и вид регрессионной модели.

Часто в распоряжении исследователя имеется дополнительная информация о параметрах модели. Правильный учет такой информации позволяет построить линейные оценки, обладающие лучшими свойствами, чем МНК-оценки.

Простейшим видом априорной информации о параметрах является вид линейных ограничений типа равенств. В этом случае получающаяся задача оценки параметров легко приводится к стандартной, рассматривавшейся в § 1 (см. п. 2.2).

Более сложно учитывать априорную информацию о параметрах модели, заданную в виде ограничений типа неравенств. Ниже предполагается, что указанные ограничения задают параметрическое множество Ω , являющееся эллипсоидом:

$$\Omega = \{\theta \in \mathbb{R}^n \mid (\theta - \theta_0)^T B (\theta - \theta_0) \leq k\}. \quad (2.1)$$

Здесь B — положительно определенная матрица, $\theta_0 \in \mathbb{R}^n$ — центр эллипсоида, $k > 0$. Априорная информация состоит в том, что $\theta \in \Omega$. При такой априорной информации наиболее естественный подход — минимаксный, рассмотренный в п. 2.3.

Часто дополнительная информация о параметрах модели $(Y, F\theta, \sigma^2 W)$ представима в виде

$$r = R\theta + \varphi, \quad \varphi \sim (0, V), \quad V > 0, \quad E \varphi \varphi^T = 0, \quad (2.2)$$

где r — некоторый известный l -вектор, R — фиксированная $l \times m$ -матрица, V — известная $l \times l$ -матрица. В этом случае r играет роль вектора наблюдений в линейной регрессионной модели $(r, R\theta, V)$, φ — вектор ошибок измерений, причем случайные векторы φ и ε некоррелированы. НЛН-оценки параметров θ при наличии априорной информации (2.2) построены в п. 2.4.

Априорная информация вида (2.2) встречается в том случае, когда в распоряжении исследователя имеются дополнительные выборочные данные по изучаемой или аналогичной модели и когда уже имелась предварительная оценка $\hat{\theta}$ всех неизвестных параметров (в этом случае $r = \hat{\theta}$, $R = I_m$, $l = m$) или их части. Отметим, что априорная информация, рассматриваемая в п. 2.2, может быть записана в виде (2.2). Действительно, если $V = 0$, то (2.2) сводится к линейным ограничениям $R\theta = r$.

В п. 2.5 предполагается, что на σ -алгебре подмножеств множества Ω имеется априорное распределение с заданным вектором средних и дисперсионной матрицей. Как и в п. 2.4, задача нахождения НЛН-оценок в этом случае сводится к классической.

2.2. Линейные ограничения на параметры. Пусть в рамках схемы $(Y, F\theta, \sigma^2 W)$ имеются ограничения на параметры θ вида $R\theta = r$, где R есть $q \times m$ -матрица ранга q ($0 < q < m$), r — q -вектор. Согласно следствию 1.1 из приложения 1 общее решение системы уравнений $R\theta = r$ имеет вид $\theta = \theta_0 + B\beta$, где θ_0 — частное решение, β — произвольный $(m - q)$ -вектор, а B есть $m \times (m - q)$ -матрица ранга $m - q$. Положим $Z = Y - F\theta_0$. Имеем

$$EZ = EY - F\theta_0 = F(\theta_0 + B\beta) - F\theta_0 = FB\beta.$$

Следовательно, рассматриваемая схема регрессии эквивалентна схеме $(Z, FB\beta, \sigma^2 W)$, в которой роль вектора результатов измерений играет $Z = Y - F\theta_0$, а новыми параметрами являются β . Несмещенная оценка параметров β может быть найдена в случае, когда матрица FB имеет полный ранг. Оценкой параметров θ является $\bar{\theta} = \theta_0 + B\bar{\beta}$, где $\bar{\beta}$ — оценка для β .

2.3. Минимаксное оценивание. Пусть $a \in R^m$ — фиксированный ненулевой вектор. Запишем функцию, зависящую от неизвестных параметров θ и их оценок $\bar{\theta}$, в виде

$$r(\theta, \bar{\theta}) = r(\theta, \bar{\theta}, a) = a^T [E(\bar{\theta} - \theta)(\bar{\theta} - \theta)^T] a. \quad (2.3)$$

Качество оценки в данном пункте будем измерять величиной

$$\max_{\theta \in \Omega} r(\theta, \bar{\theta}), \quad (2.4)$$

а линейную по измерениям оценку

$$\bar{\theta} = \arg \min_{\bar{\theta} \in \Omega} [\max_{\theta \in \Omega} r(\theta, \bar{\theta})], \quad (2.5)$$

на которой достигается минимум указанной величины, будем называть *минимаксной оценкой*.

Теорема 2.1. Если имеется схема регрессии $(Y, F\theta, \sigma^2 W)$ и параметрическое множество Ω имеет вид (2.1), то минимаксная оценка (2.5) может быть записана в виде

$$\bar{\theta}_{(M)} = \left[\frac{\sigma^2}{k} B + F^T W^{-1} F \right]^{-1} F^T W^{-1} (Y - F\theta_0) + \theta_0. \quad (2.6)$$

Доказательство. Не ограничивая общности, положим $\theta_0 = 0$ (если $\theta_0 \neq 0$, то заменяем θ на $\theta - \theta_0$ и делаем соответствующие изменения в формулах для оценок).

Для любой линейной однородной оценки $\bar{\theta} = CY$ имеем

$$\bar{\theta} - \theta = (CF - I_m) \theta + C\varepsilon, \quad r(\theta, \bar{\theta}, a) = \sigma^2 a^T C W C^T a + \theta^T B^{1/2} \bar{a} \bar{a}^T B^{1/2} \theta, \quad (2.7)$$

где симметричная $m \times m$ -матрица $B^{1/2}$ определяется из условия $(B^{1/2})^2 = B$, $\bar{a} = (B^{1/2})^{-1} (CF - I_m)^T a$.

В силу теоремы 1.9 из приложения 1 для любой $m \times m$ -матрицы A имеем

$$\max_{\theta \neq 0} (\theta^T B^{1/2} A B^{1/2} \theta) / (\theta^T B \theta) = \lambda_{\max}(A). \quad (2.8)$$

Применяя (2.8) к (2.7) при $A = \bar{a}\bar{a}^T$, получаем

$$\max_{\theta^T B \theta \leq k} r(\theta, \bar{0}, a) = \sigma^2 a^T C W C^T a + k \lambda_{\max}(A).$$

Обозначим правую часть этого равенства через $J(C)$. Поскольку $A = \bar{a}\bar{a}^T$, то $\lambda_{\max}(A) = \bar{a}^T \bar{a}$, и поэтому

$$J(C) = \sigma^2 a^T C W C^T a + k a^T (C F - I_m) B^{-1} (C F - I_m)^T a. \quad (2.9)$$

Возьмем производную от (2.9) по C :

$$\frac{dJ}{dC} = 2[(\sigma^2 W + k F B^{-1} F^T) C^T a a^T - k F B^{-1} a a^T]$$

(здесь использованы результаты § 3 приложения 1). Приравняв эту производную нулю, получаем, что минимальное значение функционала $J(C)$ достигается на матрице

$$C_* = k B^{-1} F^T (\sigma^2 W + k F B^{-1} F^T)^{-1}.$$

С помощью элементарных преобразований получаем

$$(\sigma^2 B + k F^T W^{-1} F) C_* = k \sigma^2 F^T W^{-1} W (\sigma^2 W + k F^T W^{-1} F)^{-1} + \\ + k F^T W^{-1} (k F B^{-1} F^T) (\sigma^2 W + k F^T W^{-1} F)^{-1} = k F^T W^{-1}.$$

Отсюда следует (2.6) при $\theta_0 = \theta$. Теорема доказана.

Важным свойством минимаксной оценки (2.6) является то, что ее вид не зависит от вектора a , фигурирующего в (2.3).

В частном случае $\theta_0 = 0$ и $W = I_N$ (т. е. центр эллипсоида (2.1) находится в начале координат и линейная регрессионная модель является классической) минимаксная оценка (2.6) принимает вид

$$\bar{\theta}_{(M)} = \left[\frac{\sigma^2}{k} B + F^T F \right]^{-1} F^T Y. \quad (2.10)$$

Сравнивая (2.10) и (1.59), видим, что гребневая оценка является минимаксной линейной при наличии априорной информации вида $\theta^T \theta \leq \sigma^2/k$.

При больших k (т. е. при малой априорной информации) оценка (2.6) близка к обобщенной МНК-оценке (1.48) и в пределе (при $k \rightarrow \infty$) совпадает с ней. Отсюда следует, что МНК-оценки являются минимаксными для случая, когда априорная информация о параметрах отсутствует.

Теорема 2.2. В классе линейных несмещенных оценок обобщенная МНК-оценка (1.48) является минимаксной.

Доказательство. Из условия несмещенности оценки $\bar{\theta} = C Y$ имеем $C F = I_m$, и поэтому в (2.7) $\bar{a} = 0$, в силу чего

$$\max_{\theta^T \theta \leq k} r(\bar{\theta}, \theta, a) = \sigma^2 a^T C W C^T a = a^T D \bar{\theta} a.$$

В силу результатов п. 1.7 минимум в последнем выражении для любого $a \in R^m$ достигается при $C = (F^T W^{-1} F)^{-1} F^T W^{-1}$. Теорема доказана.

2.4. Смешанное оценивание. Смешанной моделью называется модель $(Y, F\theta, \sigma^2 W)$ при наличии априорной информации (2.2). Смешанная модель может быть записана в виде

$$(\bar{Y}, \bar{F}\theta, \sigma^2 \bar{W}),$$

где

$$\bar{Y} = \begin{bmatrix} Y \\ r \end{bmatrix}, \quad \bar{F} = \begin{bmatrix} F \\ R \end{bmatrix}, \quad \bar{W} = \begin{bmatrix} W & 0 \\ 0 & \sigma^{-2} V \end{bmatrix}.$$

Применяя результаты п. 1.7 для модели $(\bar{Y}, \bar{F}\theta, \sigma^2 \bar{W})$, получаем, что оценка

$$\begin{aligned} \hat{\theta}(\sigma^2) &= (\bar{F}^T \bar{W}^{-1} \bar{F})^{-1} \bar{F}^T \bar{W}^{-1} \bar{Y} = \\ &= (\sigma^{-2} F^T W^{-1} F + R^T V^{-1} R)^{-1} (\sigma^{-2} F^T W^{-1} Y + R^T V^{-1} r) \end{aligned} \quad (2.11)$$

является НЛН-оценкой для смешанной модели. Поскольку значение σ^2 обычно неизвестно, также неизвестна и $\hat{\theta}(\sigma^2)$. В том случае, когда дисперсии и ковариации случайного вектора ошибок ϵ пропорциональны соответствующим величинам вектора ϵ , т. е. когда $V = (\sigma^2/k) W$ (где k — известное число), оценка (2.11) принимает вид

$$\hat{\theta} = (F^T W^{-1} F + k R^T V^{-1} R)^{-1} (F^T W^{-1} Y + k R^T V^{-1} r). \quad (2.12)$$

В общем случае на практике ограничиваются подстановкой в (2.11) вместо σ^2 какой-либо оценки для σ^2 (например, s^2 для смешанной модели).

2.5. Байесовское оценивание. Пусть на σ -алгебре подмножеств множества Ω задано распределение $P(d\theta)$ со средним r и дисперсионной матрицей V , независимое от распределения ϵ

$$r = \int_{\Omega} \theta P(d\theta), \quad V = \int_{\Omega} (\theta - r)(\theta - r)^T P(d\theta).$$

и пусть это распределение отражает априорные сведения о значении θ : $P(d\theta)$ является распределением вероятностей для значений неизвестных параметров. В этом случае априорная информация может быть записана в виде $(r, I_m \theta, V)$, и, следовательно, НЛН-оценкой является

$$\hat{\theta}(\sigma^2) = (\sigma^{-2} F^T W^{-1} F + V^{-1})^{-1} (\sigma^{-2} F^T W^{-1} Y + V^{-1} r). \quad (2.13)$$

Оценка (2.13) называется *байесовской оценкой*.

Отметим, что в случае $W = I_n$, $V = I_m$, $r = 0$ оценка (2.13) совпадает с гребневой оценкой (1.59). Следовательно, гребневая оценка (1.59) является байесовской для классической линейной регрессионной модели при наличии априорного распределения с известным вектором средних $r = 0$ и дисперсионной матрицей $V = I_m$.

Упражнения.

1. Положим

$$G = \left[\frac{\sigma^2}{k} B + F^T W^{-1} F \right]^{-1}.$$

Покажите, что если выполнены предположения теоремы 2.1, то для оценки (2.6) имеют место следующие равенства:

$$E\hat{\theta}_{(M)} - \theta = - \frac{\sigma^2}{k} G B (\theta - \theta_0), \quad D\hat{\theta}_{(M)} = \sigma^2 G F^T W^{-1} F G,$$

$$\sup_{\theta \in \Omega} r(\hat{\theta}_{(M)}, \theta, a) = \sigma^2 a^T G a.$$

2. Предельным переходом при $k \rightarrow \infty$ получите из (2.12) явный вид НЛН-оценки параметров θ в схеме регрессии $(Y, F\theta, \sigma^2 W)$ при наличии линейных ограничений $R\theta = 0$:

$$\bar{\theta} = (F^T W^{-1} F + R^T R)^{-1} F^T W^{-1} Y.$$

§ 3. Основные понятия дисперсионного анализа

3.1. F -критерий для проверки линейных гипотез о параметрах линейной регрессии. Предположим, что имеется схема линейной регрессии $(Y, F\theta, \sigma^2 I_N)$, в которой вектор результатов измерений Y нормально распределен:

$$Y \sim N(F\theta, \sigma^2 I_N), \quad (3.1)$$

а матрица $F^T F$ не обязательно имеет полный ранг. Пусть T — матрица размера $k \times m$, $\text{rang } T = k < m$, и строки матрицы T — линейные комбинации строк матрицы F , т. е. существует такая $k \times N$ -матрица L , что $T = LF$. При этом условии, как следует из леммы 1.6, k параметрических функций $t_i \theta$ ($i = 1, \dots, k$) оцениваемы. Здесь $t_1, \dots, t_k$ — строки матрицы T . Обозначим МНК-оценку векторной параметрической функции $\tau = T\theta$ через $\hat{\tau}$.

В силу теоремы 1.5 $\hat{\tau} = T\bar{\theta}$, где $\bar{\theta}$ — любое решение системы нормальных уравнений $F^T F \theta = F^T Y$, т. е. $\bar{\theta}$ есть МНК-оценка параметров θ , записываемая в виде (см. 1.56))

$$\bar{\theta} = (F^T F)^- F^T Y, \quad (3.2)$$

где матрица $(F^T F)^-$ (обобщенная обратная для $F^T F$) обладает тем свойством, что

$$F^T F (F^T F)^- F^T F = F^T F. \quad (3.3)$$

По лемме 1.5 величина

$$R_0^2 = (Y - F\bar{\theta})^T (Y - F\bar{\theta}), \quad (3.4)$$

которая называется *остаточной суммой квадратов*, определяется однозначно. То же самое относится и к оценке

$$\hat{\tau} = T\bar{\theta}. \quad (3.5)$$

Покажем, что величины R_0^2 и $\hat{\tau}$ независимы.

Теорема 3.1. Для схемы нормальной регрессии (3.1) в случае $T = LF$ статистики (3.4) и (3.5) независимы.

Доказательство. Преобразуем выражения (3.4) и (3.5):

$$\begin{aligned} R_0^2 &= (Y - F\hat{\theta})^T (Y - F\hat{\theta}) = Y^T Y - Y^T F\hat{\theta} - \hat{\theta}^T F^T Y + \hat{\theta}^T F^T F\hat{\theta} = \\ &= Y^T Y - Y^T F (F^T F)^{-1} F^T Y - Y^T F (F^T F)^{-1} F^T Y + \\ &+ Y^T F (F^T F)^{-1} F^T F (F^T F)^{-1} F^T Y = Y^T Y - Y^T F (F^T F)^{-1} F^T Y = \\ &= Y^T (I_N - F (F^T F)^{-1} F^T) Y, \\ \hat{\tau} &= T\hat{\theta} = LF (F^T F)^{-1} F^T Y. \end{aligned}$$

Следовательно,

$$R_0^2 = Y^T A Y, \quad \hat{\tau} = B Y, \quad (3.6)$$

где

$$A = I_N - F (F^T F)^{-1} F^T, \quad B = LF (F^T F)^{-1} F^T. \quad (3.7)$$

Покажем, что $BA = 0$. Имеем $BA = LG (F^T F)^{-1} F^T$, где

$$G = F - F (F^T F)^{-1} F^T F. \quad (3.8)$$

Подсчитаем $G^T G$:

$$\begin{aligned} G^T G &= (F^T - F^T F (F^T F)^{-1} F^T) (F - F (F^T F)^{-1} F^T F) = \\ &= F^T F - F^T F (F^T F)^{-1} F^T F - F^T F (F^T F)^{-1} F^T F + \\ &+ F^T F (F^T F)^{-1} F^T F (F^T F)^{-1} F^T F. \end{aligned}$$

Из (3.3) вытекает, что $G^T G = 0$ и, следовательно, $G = 0$; поэтому $BA = 0$.

Теперь покажем, что из условия $BA = 0$ следует независимость квадратичной и линейной форм (3.6). Воспользуемся теоремой 1.1 из приложения 1 для спектрального разложения

матрицы A , представив ее в виде $A = \sum_{i=1}^N \lambda_i P_i P_i^T$, где λ_i и P_i — собственные числа и ортонормированные собственные векторы матрицы A . Теперь из $BA = 0$ следует, что $B P_i = 0$ ($i = 1, \dots, N$). Поэтому вектор $B Y$ некоррелирован со случайными величинами $P_i^T Y$ ($i = 1, \dots, N$) и в силу нормальности Y не зависит от них. Утверждение теоремы следует из того, что

$$R_0^2 = Y^T A Y = \sum_{i=1}^N \lambda_i (P_i^T Y)^2.$$

Теорема доказана.

Лемма 3.1. Пусть имеется схема регрессии (3.1). Тогда

$$\hat{\tau} \sim N(T\theta, \sigma^2 T (F^T F)^{-1} T^T), \quad (3.9)$$

где $\hat{\tau}$ — МНК-оценка векторной параметрической функции $\tau = T\theta$, T — матрица, представимая в виде $T = LF$.

Доказательство. Поскольку случайный вектор Y имеет нормальное распределение, то случайный вектор $\hat{\tau}$ также нормально распределен. Имеем

$$E\hat{\tau} = E T \hat{\theta} = E T (F^T F)^{-1} F^T Y = T (F^T F)^{-1} F^T F \theta = L F (F^T F)^{-1} F^T F \theta.$$

При доказательстве теоремы 3.1 было показано, что матрица (3.8) состоит из нулей. Следовательно,

$$F = F (F^T F)^{-1} F^T F, \quad (3.10)$$

откуда $E\hat{\tau} = T\theta$.

Дисперсионную матрицу вектора $\hat{\tau}$ вычислим, используя (3.10) и лемму 1.3:

$$\begin{aligned} D\hat{\tau} &= D L F \hat{\theta} = D L F (F^T F)^{-1} F^T Y = L F (F^T F)^{-1} F^T [D Y] F (F^T F)^{-1} \times \\ &\times F^T L^T = \sigma^2 L F (F^T F)^{-1} F^T F (F^T F)^{-1} F^T L^T = \sigma^2 L F (F^T F)^{-1} F^T L^T = \\ &= \sigma^2 T (F^T F)^{-1} T^T. \end{aligned}$$

Следовательно,

$$D\hat{\tau} = \sigma^2 V, \text{ где } V = T (F^T F)^{-1} T^T. \quad (3.11)$$

Лемма доказана.

Лемма 3.2. Для схемы регрессии (3.1), в которой $N \times m$ -матрица F имеет ранг r ($\text{rang } F = r \leq \min(m, N)$) остаточная сумма квадратов (3.4), умноженная на σ^2 , имеет χ^2 -распределение с $N - r$ степенями свободы:

$$R_0^2 \sim \sigma^2 \chi^2(N - r). \quad (3.12)$$

Доказательство. Используя представление (3.6) для R_0^2 и теорему 1 из приложения 2, получаем

$$R_0^2 \sim \sigma^2 \chi^2(\text{tr } A),$$

где матрица A определяется по (3.7).

Осталось показать, что $\text{tr } A = N - r$. Имеем

$$\text{tr } A = \text{tr}(I_N - F (F^T F)^{-1} F^T) = \text{tr } I_N - \text{tr } F (F^T F)^{-1} F^T.$$

По лемме 1.5 матрица $F (F^T F)^{-1} F^T$ определена однозначно, т. е. не зависит от выбора способа обобщенного обращения матрицы. Выберем способ, определяемый формулой (1.15) из приложения 1. Тогда

$$\text{tr } F (F^T F)^{-1} F^T = \text{tr } F^T F (F^T F)^{-1} = \text{rang } F^T F = \text{rang } F = r. \quad (3.13)$$

Лемма доказана.

Задача проверки общих линейных гипотез о параметрах линейной регрессии ставится следующим образом. Пусть имеется схема регрессии (3.1), где $N \times m$ -матрица F имеет ранг $r \leq \min(m, N)$, $k \times m$ -матрица T имеет ранг k ($k \leq r$) и допускает представление $T = L F$. Общую линейную гипотезу о параметрах θ модели (3.1) запишем в виде

$$H_0: T\theta = \tau_0, \quad (3.14)$$

где τ_0 — фиксированный k -вектор.

Для МНК-оценки $\hat{\tau}$ векторной параметрической функции $\tau = T\theta$ выполнено (3.9). Используя формулу (4) из приложения 2, получаем

$$(\hat{\tau} - \tau)^T V^{-1} (\hat{\tau} - \tau) \sim \sigma^2 \chi^2(k), \quad (3.15)$$

где матрица V определена в (3.11).

Из (3.12) и (3.15) следует, что при справедливости гипотезы (3.14) выполнено

$$\mathcal{F} = \frac{(\hat{\tau} - \tau_0)^T V^{-1} (\hat{\tau} - \tau_0)}{k} / \frac{R_0^2}{N-r} \sim F(k, N-r), \quad (3.16)$$

т. е. статистика $\mathcal{F}$ имеет распределение Фишера с k и $N-r$ степенями свободы (см. приложение 2).

Теперь выясним тенденцию поведения статистики $\mathcal{F}$ в случае, когда гипотеза H_0 неверна. Используя (3.9), имеем

$$\begin{aligned} \frac{E(\hat{\tau} - \tau_0)^T V^{-1} (\hat{\tau} - \tau_0)}{k} &= \\ &= \frac{1}{k} \text{tr} V^{-1} E(\hat{\tau} - \tau + \tau - \tau_0)(\hat{\tau} - \tau + \tau - \tau_0)^T = \\ &= \frac{1}{k} \text{tr} V^{-1} [E(\hat{\tau} - \tau)(\hat{\tau} - \tau)^T + (\tau - \tau_0)(\tau - \tau_0)^T] = \\ &= \frac{1}{k} \text{tr} V^{-1} [\sigma^2 V + (\tau - \tau_0)(\tau - \tau_0)^T] = \\ &= \frac{\sigma^2}{k} \text{tr} I_k + \frac{1}{k} (\tau - \tau_0)^T V^{-1} (\tau - \tau_0) = \\ &= \sigma^2 + \frac{1}{k} (\tau - \tau_0)^T V^{-1} (\tau - \tau_0). \end{aligned}$$

Следовательно, равенство

$$E \frac{(\hat{\tau} - \tau_0)^T V^{-1} (\hat{\tau} - \tau_0)}{k} = \sigma^2$$

имеет место только в случае справедливости гипотезы H_0 .

Поскольку $(\tau - \tau_0)^T V^{-1} (\tau - \tau_0) > 0$ при всех $\tau \neq \tau_0$, то при невыполнении гипотезы H_0 выполняется неравенство

$$E \frac{(\hat{\tau} - \tau_0)^T V^{-1} (\hat{\tau} - \tau_0)}{k} > \sigma^2, \quad (3.17)$$

причем, чем больше τ отличается от τ_0 , тем больше разница между правой и левой частью (3.17). С другой стороны, безотносительно к справедливости гипотезы H_0 выполнено

$$E \frac{R_0^2}{N-r} = \sigma^2 \quad (3.18)$$

(вывод этой формулы аналогичен доказательству теоремы 1.2 и использует (3.13)).

Таким образом, при невыполнении гипотезы H_0 числитель правой части (3.16) в среднем превосходит знаменатель, причем

величина этого превышения зависит от того, насколько τ_0 отличается от истинного значения τ . Следовательно, большие значения статистики $\mathcal{F}$ свидетельствуют о том, что выборочные данные противоречат гипотезе H_0 .

Если выбрать величину F_α из условия

$$P\{F > F_\alpha\} = \alpha, \quad \text{где } F \sim F(k, N-r), \quad (3.19)$$

то при справедливости гипотезы H_0 вероятность выполнения неравенства

$$\mathcal{F} = \frac{(\hat{\tau} - \tau_0)^T V^{-1} (\hat{\tau} - \tau_0)}{k} / \frac{R_1^2}{N-r} > F_\alpha \quad (3.20)$$

будет равна α . Поэтому, если α мало и выполнено неравенство, противоположное (3.20), можно считать, что выборочные данные не противоречат гипотезе H_0 (т. е. гипотезу H_0 следует принять). Если же выполнено (3.20), то гипотезу H_0 следует отвергнуть, поскольку она не согласуется с результатами измерений.

В частном случае при $k=1$, когда проверяется гипотеза только об одном параметре регрессии, описанная процедура проверки гипотезы H_0 эквивалентна процедуре, описанной в п. 1.6.

Для вычисления статистики $\mathcal{F}$ из левой части (3.20) можно пользоваться другой процедурой, которая не требует явного вычисления МНК-оценок векторной параметрической функции $\tau = T\theta$. Эта процедура основана на использовании следующей леммы.

Лемма 3.3. Пусть имеются схема регрессии $(Y, F\theta, \sigma^2 I_N)$ и $k \times m$ -матрица T , допускающая представление $T = LF$. Тогда

$$(\hat{\tau} - \tau)^T V^{-1} (\hat{\tau} - \tau) = R_1^2 - R_0^2, \quad (3.21)$$

где $\hat{\tau}$ — МНК-оценка векторной параметрической функции $T\theta$, матрица V определена в (3.11) и

$$R_1^2 = \min_{\theta: T\theta = \tau} (Y - F\theta)^T (Y - F\theta) \quad (3.22)$$

(τ_0 — фиксированный k -вектор).

Доказательство. С помощью метода множителей Лагранжа (см. [2]) получаем, что минимум в (3.22) достигается на векторе θ_* , являющемся при некотором $\lambda = \lambda_*$ решением системы уравнений

$$\begin{aligned} F^T F \theta + T^T \lambda &= F^T Y, \\ T \theta &= \tau_0. \end{aligned} \quad (3.23)$$

Поскольку

$$\begin{aligned} R_1^2 &= (Y - F\theta_*)^T (Y - F\theta_*) = \\ &= (Y - F\hat{\theta})^T (Y - F\hat{\theta}) + (\hat{\theta} - \theta_*)^T F^T F (\hat{\theta} - \theta_*), \end{aligned}$$

где $\hat{\theta} = (F^T F)^{-1} F^T Y$ — оценка МНК параметров θ , то для доказательства справедливости (3.21) необходимо показать, что

$$(\hat{\theta} - \theta_*)^T F^T F (\hat{\theta} - \theta_*) = (\hat{\tau} - \tau_0)^T V^{-1} (\hat{\tau} - \tau_0). \quad (3.24)$$

Матрица T допускает представление $T = LF$. Отсюда следует, что она допускает и представление $T = CF^T F$, где C — некоторая $k \times m$ -матрица. Из (3.23) находим

$$F^T F (\hat{\theta} - \theta_0) = T^T \lambda_0 = F^T C F^T \lambda_0.$$

Поэтому

$$(\hat{\theta} - \theta_0)^T F^T F (\hat{\theta} - \theta_0) = \lambda_0^T C F^T C F^T \lambda_0. \quad (3.25)$$

Далее из (3.23) получаем

$$C F^T F \hat{\theta}_0 - T \hat{\theta}_0 + C T^T \lambda_0 = C F^T Y - \tau_0,$$

откуда

$$C T^T \lambda_0 = \hat{t} - \tau_0.$$

Это равенство переписывается в виде

$$C F^T C F^T \lambda_0 = \hat{t} - \tau_0.$$

Отсюда

$$\lambda_0 = V^{-1}(\hat{t} - \tau_0).$$

Подставляя это выражение в (3.25), получаем (3.24). Лемма доказана.

Из (3.21) вытекает, что статистику F , находящуюся в левой части (3.20), можно переписать в виде

$$\mathcal{F} = \frac{R_1^2 - R_0^2}{k} / \frac{R_0^2}{N - r}. \quad (3.26)$$

Вычисления, приводящие к статистике $\mathcal{F}$, часто располагают в виде так называемой *таблицы дисперсионного анализа*.

Таблица 3.1

	Число степеней свободы	Сумма квадратов	Средний квадрат
Остаточная сумма	$N - r$	$R_0^2 = \min_{\theta} (Y - F\theta)^T (Y - F\theta)$	$\frac{R_0^2}{N - r}$
Полная сумма	$N - r + k$	$R_1^2 = \min_{\theta: T\theta = \tau_0} (Y - F\theta)^T (Y - F\theta)$	
Отклонение от гипотезы	k	$R_1^2 - R_0^2$	$\frac{R_1^2 - R_0^2}{N - r}$

Задачи проверки линейных гипотез вида (3.14) о параметрах линейной регрессионной модели (3.1) называют *задачами дисперсионного анализа*. При этом часто предполагают, что переменные, от которых зависит функция регрессии, являются качественными, т. е. принимают конечное число значений, соответствующих каким-либо качествам. В задачах дисперсионного анализа указанные переменные принято называть *факторами*. Значения, которые принимает тот или иной фактор, называют его *уровнями*. Множество допустимых уровней количественного

фактора обычно представляется в виде отрезка прямой, качественного — в виде нескольких чисел.

Часто качественный фактор принимает значения на двух уровнях 0 и 1, символизирующих «нет» и «да». Это может быть, например, отсутствие или наличие того или иного ингредиента в смеси, неучастие или участие того или иного вещества в реакции, отсутствие или присутствие при проведении эксперимента некоторого лица и т. п. В других случаях число уровней качественного фактора l больше двух (например, l — число различных сортов семян в сельскохозяйственном эксперименте или различных доз лекарства в медицинском).

Возникающие в дисперсионном анализе задачи планирования эксперимента носят специфический характер и направлены в основном на уменьшение числа проводимых измерений и на упрощение процедуры анализа. Эти задачи рассмотрены в гл. 2, а в п. 3.2 в качестве примера использования изложенной выше общей методики проверки гипотез рассмотрим задачу однофакторного дисперсионного анализа, являющуюся одной из простейших задач подобного рода.

3.2. Однофакторный дисперсионный анализ. Предположим, что имеется m повторных выборок

$$Y_1 = (y_{11}, \dots, y_{1N_1})^T, \dots, Y_m = (y_{m1}, \dots, y_{mN_m})^T \quad (3.27)$$

объемов $N_1, \dots, N_m$ из значений нормально распределенных случайных величин с неизвестными средними $\theta_1, \dots, \theta_m$ и неизвестной общей дисперсией σ^2 . Предполагается, что измерения $y_{ij} = \theta_i + \varepsilon_{ij}$ ($i = 1, \dots, m$; $j = 1, \dots, N_i$) проводились при разных значениях некоторого фактора, влияние которого может сказываться на значениях среднего θ_i . Гипотеза

$$H_0: \theta_1 = \theta_2 = \dots = \theta_m \quad (3.28)$$

означает, что указанные факторы на результаты измерений не влияют.

Задача проверки гипотезы (3.28) по результатам измерений (3.27) может быть записана в виде задачи проверки гипотезы (3.14) о параметрах регрессионной модели (3.1), если положить $N = N_1 + \dots + N_m$, $\theta = (\theta_1, \dots, \theta_m)^T$,

$$Y = (y_{11}, \dots, y_{1N_1}, y_{21}, \dots, y_{2N_2}, \dots, y_{m1}, \dots, y_{mN_m})^T,$$

$$F^T = \begin{bmatrix} 1 & 1 & \dots & 1 & 0 & 0 & \dots & 0 & 0 & \dots & 0 \\ 0 & 0 & \dots & 0 & 1 & 1 & \dots & 1 & 0 & 0 & \dots & 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & \dots & 0 & 0 & 0 & \dots & 0 & \dots & 1 & 1 & \dots & 1 \end{bmatrix},$$

$\underbrace{\hspace{1.5cm}}_{N_1} \quad \underbrace{\hspace{1.5cm}}_{N_2} \quad \underbrace{\hspace{1.5cm}}_{N_m}$

$$B = \begin{bmatrix} 1 & 0 & \dots & 0 & -1 \\ 0 & 1 & \dots & 0 & -1 \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & \dots & 1 & -1 \end{bmatrix},$$

$$\tau_0 = 0, \quad k = m - 1, \quad r = m.$$

Если гипотеза (3.28) верна, т. е. все θ_i ($i = 1, \dots, m$) равны θ_1 , то минимум суммы квадратов отклонений $\sum_{i, l} (y_{il} - \theta_i)^2$ равен

$$R_1^2 = \min_{\theta_1} \sum_{i, l} (y_{il} - \theta_1)^2 = \sum_{i, l} y_{il}^2 - \frac{1}{N} \left(\sum_{i, l} y_{il} \right)^2. \quad (3.29)$$

Остаточная сумма квадратов R_0^2 равна

$$\begin{aligned} R_0^2 &= \min_{\theta_1, \dots, \theta_m} \sum_{i, l} (y_{il} - \theta_i)^2 = \\ &= \sum_i \min_{\theta_i} \sum_l (y_{il} - \theta_i)^2 = \sum_{i=1}^m \left(\sum_{l=1}^{N_i} y_{il}^2 - B_i^2 / N_i \right), \end{aligned}$$

где

$$B_i = \sum_{l=1}^{N_i} y_{il}.$$

Следовательно, сумма квадратов за счет отклонения от гипотезы (3.28) имеет $m - 1$ степеней свободы и равна

$$R_1^2 - R_0^2 = \frac{B_1^2}{N_1} + \dots + \frac{B_m^2}{N_m} - \frac{B^2}{N}, \quad (3.30)$$

где

$$B = B_1 + \dots + B_m = \sum_{i, l} y_{il}.$$

На практике проще вычислять R_1^2 и $R_1^2 - R_0^2$ по формулам (3.29) и (3.30), а затем вычислять R_0^2 . Соответствующая таблица дисперсионного анализа выглядит следующим образом.

Т а б л и ц а 3.2. Однофакторный дисперсионный анализ

	Число степеней свободы	Сумма квадратов	Средний квадрат
Остаточная сумма (расхождение внутри выборок)	$N - m$	R_0^2 (получается вычитанием)	$\frac{R_0^2}{N - m}$
Полная сумма	$N - 1$	$R_1^2 = \sum_{i, l} y_{il}^2 - \frac{B^2}{N}$	
Отклонение от гипотезы (расхождение между выборками)	$m - 1$	$R_1^2 - R_0^2 = \sum_{i=1}^m \frac{B_i^2}{N_i} - \frac{B^2}{N}$	$\frac{R_1^2 - R_0^2}{m - 1}$

Упражнения.

1. Пусть модель регрессии и результаты измерений те же, что и в примере 1.7, и пусть значение σ^2 неизвестно. Проверьте гипотезы: а) $\theta_1 = \theta_2$; б) $\theta_1 = \theta_2/2$; в) $\theta_1 = \theta_2 + 1$; г) $\theta_1 = \theta_2 - 1$.

2. Пусть схема регрессии та же, что и в примере 1.10. Сформулируйте задачу однофакторного дисперсионного анализа. Проведите однофакторный дисперсионный анализ для данных примера 1.10.

Глава 2

ФАКТОРНЫЕ ПЛАНЫ

Исторически теория планирования эксперимента начала развиваться с факторного планирования. Основы факторного планирования были заложены еще в 30-х годах нашего столетия. Суть теории факторного планирования состоит в построении экономических планов, по результатам измерений в точках которых можно проводить просто реализуемые процедуры статистических выводов о неизвестных параметрах полиномиальных функций регрессии. Эта теория широко используется на практике и неразрывно связана с такими областями комбинаторики, как блок-схемы, латинские квадраты, конечные геометрии.

В данной главе дается элементарное введение в факторное планирование. Описаны простейшие методы построения дробных факторных планов, ортогональных латинских квадратов и сбалансированных блок-схем, процедура проведения дисперсионного анализа для латинских планов и соотношения между параметрами сбалансированных неполных блок-схем.

§ 1. Полные факторные планы и их дробные реплики

1.1. Понятие о факторном планировании. Рассмотрим схему регрессионного эксперимента, в которой результаты измерений зависят от значений некоторого числа $m > 1$ переменных x_i ($i = 1, \dots, m$), называемых *факторами*. При этом предполагается, что фактор x_i может принимать конечное число значений $s_i \geq 2$, которые называются *уровнями фактора*, и что выбор уровней факторов находится в распоряжении экспериментатора.

Метод, состоящий в рассмотрении влияния факторов на результаты измерений по одному, называется *классическим экспериментом*. В отличие от него в *методе факторного планирования* уровни всех факторов комбинируются. Как правило, при равном числе измерений оценки неизвестных параметров регрессии, получающиеся для разумно спланированного факторного эксперимента, более точны.

Схема регрессии, в рамках которой проводится тот или иной факторный эксперимент, называется *факторной моделью*. Факторные модели могут быть различными. В качестве функции

регрессии в этих моделях при $s_i \equiv 2$ обычно рассматривается полином степени $k \leq m$, в котором присутствуют только члены вида

$$x_1^{r_1} x_2^{r_2} \dots x_m^{r_m}, \quad (1.1)$$

где все r_i равны либо нулю, либо единице. В общем случае такая функция регрессии имеет вид

$$\eta(x_1, \dots, x_m) = \theta_0 + \sum_{i=1}^m \theta_i x_i + \sum_{i < j} \theta_{ij} x_i x_j + \\ + \sum_{i < j < l} \theta_{ijl} x_i x_j x_l + \dots + \sum_{i_1 < i_2 < \dots < i_m} \theta_{i_1 i_2 \dots i_m} x_{i_1} x_{i_2} \dots x_{i_m}, \quad (1.2)$$

где некоторые из параметров $\theta_i, \theta_{ij}, \dots$ неизвестны, а некоторые априори равны нулю. Параметр θ_0 в (1.2) называется *общим средним*, параметры θ_i ($i = 1, \dots, m$) — *главными эффектами* (взаимодействиями нулевого порядка), θ_{ij} — *эффектами взаимодействий первого порядка* (эффектами двухфакторных взаимодействий), θ_{ijk} — *эффектами взаимодействий второго порядка* (эффектами трехфакторных взаимодействий) и аналогично $\theta_{i_1 i_2 \dots i_k}$ — *эффектами взаимодействий порядка $k-1$* (эффектами k -факторных взаимодействий). Число k -факторных взаимодействий равно

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

В частности, имеется m главных эффектов, $m(m-1)/2$ эффектов двухфакторных взаимодействий и один эффект m -факторного взаимодействия.

Наиболее часто используется два частных случая функции регрессии (1.2):

$$\eta(x_1, \dots, x_m) = \theta_0 + \theta_1 x_1 + \dots + \theta_m x_m \text{ (линейная)}, \quad (1.3)$$

и

$$\eta(x_1, \dots, x_m) = \theta_0 + \sum_{i=1}^m \theta_i x_i + \sum_{i < j} \theta_{ij} x_i x_j \text{ (квадратичная)}; \quad (1.4)$$

последняя называется также *неполной квадратичной* в силу предположения о том, что коэффициенты θ_{ij} при одночленах x_i^2 равны нулю.

Факторный план называется *полным*, если, согласно этому плану, измерения проводятся по одному для каждой возможной комбинации уровней факторов. Эксперимент, проведенный по полному факторному плану, называется *полным факторным экспериментом*. Полный факторный план требует проведения $s_1 \dots s_m$ измерений, где s_i — число уровней фактора x_i . Если $s_i = s$ ($i = 1, \dots, m$), т. е. количество уровней всех факторов одинаково и равно s , то полный факторный план обозначается так: план s^m . Факторный план называется *дробным*, если он

предназначен для проведения числа измерений, меньшего чем s_1 .

Большое распространение на практике получили *двухуровневые планы*, т. е. такие планы, факторы в которых принимают значения только на двух уровнях. Это связано с тремя обстоятельствами. Во-первых, если какой-либо фактор x_i принимает значения на s уровнях, то можно заменить этот фактор на s новых двухуровневых факторов, соответствующих уровням фактора (т. е. рассматривая каждый уровень фактора x_i в качестве нового фактора). Во-вторых, даже если фактор может принимать значения на многих уровнях (возможно, бесконечном числе), то часто выбирают только два крайних, считая при этом, что, чем больше различие в значениях факторов, тем точнее МНК-оценки неизвестных параметров функции регрессии (если рассматривается функция регрессии (1.3), то это всегда верно). В-третьих, при проведении реального эксперимента типичной является ситуация, в которой факторы могут либо присутствовать, либо отсутствовать, т. е. имеют только два значения.

При рассмотрении двухуровневых планов уровни факторов кодируются числами 1 и -1 и условно называются соответственно *верхним* и *нижним*. Отметим, что если используется двухуровневый план, то включать в функцию регрессии слагаемые вида (1.1) с каким-либо $r_i \geq 2$ бессмысленно: действительно, если r_i четно, то $x_i^{r_i} = 1$, а если нечетно, то $x_i^{r_i} = x_i$.

Рассмотрим способы построения дробных двухуровневых планов.

1.2. Дробные реплики плана 2^m . Полный факторный план 2^m является удобным с точки зрения простоты проведения статистического анализа параметров функции регрессии (1.2). Кроме того, для линейной модели (1.3) и куба как множества планирования этот план является в определенном смысле оптимальным (см. § 1 гл. 4). Тем не менее при большом числе факторов m план 2^m используется крайне редко, так как предназначен для проведения большого числа 2^m измерений. Число 2^m быстро растет с ростом m и при больших m существенно превышает число $m + 1$ неизвестных параметров линейной функции регрессии (1.3), которая наиболее часто используется на практике.

За счет потери части информации, не очень существенной при построении линейных моделей, при $m \geq 3$ количество измерений можно существенно сократить. Для этого вместо плана 2^m следует использовать описанный ниже дробный факторный план 2^{m-p} , который предназначен для проведения 2^{m-p} измерений (натуральное число p выбирается из условия $2^{m-p} \geq m + 1$).

Сначала рассмотрим линейную функцию регрессии, зависящую от трех факторов:

$$\eta(x_1, x_2, x_3) = \theta_0 + \theta_1 x_1 + \theta_2 x_2 + \theta_3 x_3. \quad (1.5)$$

Применение полного факторного плана 2^3 , предназначенного для проведения восьми измерений в точках вида

$$(x_1, x_2, x_3) = (\pm 1, \pm 1, \pm 1),$$

позволяет, как показано ниже, несмещенно оценить не только общее среднее θ_0 и главные эффекты $\theta_1, \theta_2, \theta_3$, но также и всевозможные взаимодействия первого и второго порядков, т. е. все параметры неполной кубической модели

$$\theta_0 + \theta_1 x_1 + \theta_2 x_2 + \theta_3 x_3 + \theta_{12} x_1 x_2 + \theta_{13} x_1 x_3 + \theta_{23} x_2 x_3 + \theta_{123} x_1 x_2 x_3. \quad (1.6)$$

Для оценивания параметров функции регрессии (1.5) можно построить план, предназначенный для проведения не восьми, а четырех измерений. Для построения такого плана факторы x_1 и x_2 будем варьировать, как в плане 2^2 , а в качестве уровня фактора x_3 будем выбирать значение $x_3 = x_1 x_2$. Получим план, определяемый следующей матрицей:

Номер измерения	Матрица плана F			
	x_0	x_1	x_2	x_3
1	1	1	1	1
2	1	-1	1	-1
3	1	1	-1	-1
4	1	-1	-1	1

(1.7)

Первый столбец матрицы содержит значения фиктивной переменной x_0 , которая тождественно равна единице (параметр θ_0 является коэффициентом перед переменной x_0).

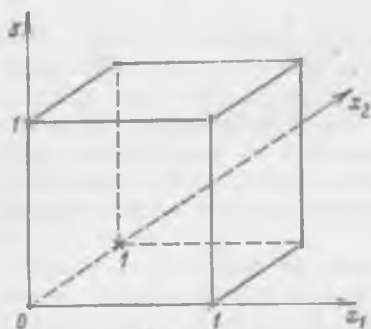


Рис. 3

Плану, определяемому матрицей (1.7), соответствуют те точки рис. 3, которые отмечены крестиками. На этом рисунке нижний уровень факторов обозначен символом θ , а все помеченные точки в совокупности определяют полный факторный план 2^3 .

По результатам измерений в точках плана, определяемого по (1.7), можно оценить параметры $\theta_0, \dots, \theta_3$ функции регрессии (1.5), так как матрица плана F невырождена.

Рассмотрим теперь функцию регрессии (1.6). Из результатов п. 1.9 гл. I вытекает, что поскольку число измерений в рассмотренном плане равно четырем, а число параметров функции

регрессии (1.6) — восьми, то все эти параметры несмещенно оценить невозможно. Для модели (1.6) матрица плана, аналогичная (1.7), выглядит следующим образом:

Номер измерения	Матрица плана F							
	x_0	x_1	x_2	x_3	x_1x_2	x_1x_3	x_2x_3	$x_1x_2x_3$
1	1	1	1	1	1	1	1	1
2	1	-1	1	-1	-1	1	-1	1
3	1	1	-1	-1	-1	-1	1	1
4	1	-1	-1	1	1	-1	-1	1

(1.8)

В выписанной матрице плана первый столбец совпадает с восьмым, второй — с седьмым, третий — с шестым, четвертый — с пятым. Следовательно, при использовании этого плана нет различия между x_0 и $x_1x_2x_3$; x_1 и x_2x_3 ; x_2 и x_1x_3 ; x_3 и x_1x_2 , т. е.

$$x_0 = x_1x_2x_3, \quad x_1 = x_2x_3, \quad x_2 = x_1x_3, \quad x_3 = x_1x_2. \quad (1.9)$$

Поэтому вместо отыскания оценок восьми параметров функции регрессии (1.6) можно найти лишь оценки четырех смешанных коэффициентов:

$$\theta_0 + \theta_{123}, \quad \theta_1 + \theta_{23}, \quad \theta_2 + \theta_{13}, \quad \theta_3 + \theta_{12}. \quad (1.10)$$

Символически это записывается так:

$$\hat{\theta}_0 \rightarrow \theta_0 + \theta_{123}, \quad \hat{\theta}_1 \rightarrow \theta_1 + \theta_{23}, \quad \hat{\theta}_2 \rightarrow \theta_2 + \theta_{13}, \quad \hat{\theta}_3 \rightarrow \theta_3 + \theta_{12}.$$

Таким образом, общее среднее и главные эффекты оцениваются независимо друг от друга, но смешиваются соответственно с эффектами взаимодействий второго и первого порядка. Если постулируется линейная модель (1.5), то эффекты взаимодействий считаются незначимыми, и набор смешанных коэффициентов (1.10) превращается в набор параметров модели (1.5).

Рассмотренный план называется *полуреplikой* или *планом* 2^{3-1} . Он получается из полного факторного плана 2^3 путем приравнивания единице произведения $x_1x_2x_3$. Действительно, из соотношения

$$1 = x_1x_2x_3 \quad (1.11)$$

следуют все четыре соотношения (1.9), при этом учитывается, что $x_0 = 1$, $x_i^2 = 1$ ($i = 1, 2, 3$). Соотношение (1.11) называется *определяющим для полуреплики* (1.8).

Полуреплика (1.8) не единственно возможная. Другая полуреплика 2^{3-1} получится, если уровни фактора x_3 устанавливать

в соответствии с равенством $x_3 = -x_1x_2$. Тогда матрица плана, аналогичная (1.8), примет вид

Номер измерения	Матрица плана F							
	x_0	x_1	x_2	x_3	x_1x_2	x_1x_3	x_2x_3	$x_1x_2x_3$
1	1	1	1	-1	1	-1	-1	-1
2	1	-1	1	1	-1	-1	1	-1
3	1	1	-1	1	-1	1	-1	-1
4	1	-1	-1	-1	1	1	1	-1

(1.12)

Определяющим соотношением для полуреплики (1.12) является

$$1 = -x_1x_2x_3. \quad (1.13)$$

Вспоминая, что $x_0 = 1$, $x_i^2 = 1$, и умножая обе части (1.13) последовательно на x_1 , x_2 , x_3 , получаем аналог соотношений (1.9):

$$x_0 = -x_1x_2x_3, \quad x_1 = -x_2x_3, \quad x_2 = -x_1x_3, \quad x_3 = -x_1x_2.$$

Отсюда следует, что при использовании полуреплики (1.12) можно несмещенно оценить четыре смешанных коэффициента:

$$\theta_0 - \theta_{123}, \quad \theta_1 - \theta_{23}, \quad \theta_2 - \theta_{13}, \quad \theta_3 - \theta_{12},$$

т. е.

$$\hat{\theta}_0 \rightarrow \theta_0 - \theta_{123}, \quad \hat{\theta}_1 \rightarrow \theta_1 - \theta_{23}, \quad \hat{\theta}_2 \rightarrow \theta_2 - \theta_{13}, \quad \hat{\theta}_3 \rightarrow \theta_3 - \theta_{12}.$$

Объединение двух полуреplik (1.8) и (1.12) дает, как нетрудно видеть, полный факторный план 2^3 .

При большом числе факторов m для оценивания параметров линейной функции регрессии (1.3) можно строить дробные реплики высокой степени дробности. Так, при $m = 7$ можно построить дробную реплику из полного факторного плана 2^3 для первых трех факторов, приравняв четыре оставшихся фактора к двухфакторным и трехфакторным взаимодействиям трех других факторов, положив, например,

$$x_4 = x_1x_2x_3, \quad x_5 = x_1x_2, \quad x_6 = x_1x_3, \quad x_7 = x_2x_3. \quad (1.14)$$

Такая реплика записывается как 2^{7-4} .

В общем случае дробную реплику будем обозначать через 2^{m-p} , если p факторов приравнены к произведениям остальных $m-p$ факторов, уровни которых выбраны согласно полному факторному плану. Дробную реплику 2^{m-p} можно строить различными способами. Для анализа системы смешивания коэффициентов пользуются понятиями генерирующих и определяющих соотношений.

Генерирующими называются соотношения, с помощью которых построена дробная реплика. Так, для реплики (1.8) генерирующим является соотношение $x_3 = x_1x_2$, для реплики (1.12) —

соотношение $x_3 = -x_1x_2$. Для указанной выше реплики 2^{7-4} генерирующими являются соотношения (1.14).

Определяющим соотношением называется равенство, в левой части которого стоит единица, а в правой — какое-либо произведение факторов. Для дробной реплики 2^{m-p} можно получить p различных определяющих соотношений из генерирующих путем умножения обеих частей последних на их левые части с последующей заменой x_i^2 на 1 ($i = 1, \dots, m$). Другие определяющие соотношения получаются путем перемножения ранее полученных и выделения среди них новых. Например, для реплики (1.8) определяющим является соотношение (1.11), для реплики (1.12) — соотношение (1.13).

Получим определяющие соотношения для реплики 2^{7-4} , задаваемой генерирующими соотношениями (1.14). Умножая обе части равенств (1.14) на их левые части, получаем четыре определяющих соотношения:

$$1 = x_1x_2x_3x_4, \quad 1 = x_1x_2x_5, \quad 1 = x_1x_3x_6, \quad 1 = x_2x_3x_7. \quad (1.15)$$

Попарное перемножение этих четырех соотношений дает шесть новых:

$$\begin{aligned} 1 &= x_3x_4x_5, \quad 1 = x_2x_4x_6, \quad 1 = x_1x_4x_7, \quad 1 = x_2x_3x_5x_6, \\ 1 &= x_1x_3x_5x_7, \quad 1 = x_1x_2x_6x_7. \end{aligned} \quad (1.16)$$

Перемножение каждой тройки из четырех соотношений (1.15) дает еще три определяющих соотношения:

$$1 = x_1x_4x_5x_6, \quad 1 = x_3x_4x_6x_7, \quad 1 = x_5x_6x_7. \quad (1.17)$$

Наконец, перемножая все четыре соотношения (1.15), получаем

$$1 = x_1x_2x_3x_4x_5x_6x_7. \quad (1.18)$$

Легко понять, что, кроме (1.15) — (1.18), других определяющих соотношений для рассмотренной реплики 2^{7-4} нет.

Знание определяющих соотношений позволяет найти всю систему совместных оценок без изучения матрицы планирования дробной реплики. Для того чтобы определить, с какими взаимодействиями смешано данное, нужно на него умножить обе части всех определяющих соотношений.

Определим, например, с какими взаимодействиями смешан главный эффект θ_3 в дробной реплике 2^{7-4} , определяемой генерирующими соотношениями (1.14). Для этого умножим все определяющие соотношения (1.15) — (1.18) на x_3 . Получим

$$\begin{aligned} x_3 &= x_1x_2x_4 = x_1x_2x_3x_5 = x_1x_6 = x_2x_7 = x_4x_5 = \\ &= x_2x_3x_4x_6 = x_1x_3x_4x_7 = x_2x_5x_6 = x_1x_5x_7 = \\ &= x_1x_2x_3x_6x_7 = x_1x_3x_4x_5x_6 = x_4x_6x_7 = x_3x_5x_6x_7 = x_1x_2x_4x_5x_6x_7. \end{aligned}$$

Следовательно, главный эффект θ_3 смешан с эффектами взаимодействий первого порядка

$$\theta_{16}, \quad \theta_{27}, \quad \theta_{45}.$$

с эффектами взаимодействий второго порядка

$$\theta_{124}, \theta_{256}, \theta_{157}, \theta_{467},$$

третьего порядка

$$\theta_{1235}, \theta_{2346}, \theta_{1347}, \theta_{3567},$$

четвертого порядка

$$\theta_{12367}, \theta_{13456}$$

и пятого порядка

$$\theta_{124567}.$$

Символически это записывается в виде

$$\begin{aligned} \hat{\theta}_3 \rightarrow \theta_3 + \theta_{16} + \theta_{27} + \theta_{45} + \theta_{124} + \theta_{256} + \theta_{157} + \theta_{467} + \theta_{1235} + \\ + \theta_{2346} + \theta_{1347} + \theta_{3567} + \theta_{12367} + \theta_{13456} + \theta_{124567} \end{aligned} \quad (1.19)$$

и означает, что МНК-оценка параметра θ_3 линейной функции регрессии (1.3) является МНК-оценкой суммы параметров, написанной в правой части (1.19), если истинная модель имеет вид (1.2) с $m = 7$.

В конкретной практической ситуации для выбора подходящей дробной реплики полного факторного плана необходимо использовать все априорные сведения теоретического и интуитивного характера об объекте планирования с целью выделения тех факторов и произведений факторов, влияние которых на результаты измерений существенно. При этом смешивание нужно производить так, чтобы общее среднее θ_0 и главные эффекты $\theta_1, \dots, \theta_m$ были смешаны с эффектами взаимодействий самого высокого порядка (так как обычно они отсутствуют) или с эффектами таких взаимодействий, о которых известно, что они оказывают несущественное влияние на результаты измерений. Отсюда следует, в частности, что недопустимо произвольное разбиение полного факторного плана 2^3 на две части для выделения полуреплики 2^{3-1} .

Качество дробного факторного плана иногда характеризуют с помощью *разрешающей способности плана*, которая равна наименьшему числу символов в правых частях определяющих соотношений. В частности, для плана разрешающей способности III ни один главный эффект не смешан ни с каким другим главным эффектом, но главные эффекты смешаны с эффектами двухфакторных взаимодействий. Для плана разрешающей способности IV главные эффекты не смешаны друг с другом и с эффектами двухфакторных взаимодействий, но последние друг с другом смешаны. Для плана разрешающей способности V главные эффекты и эффекты двухфакторных взаимодействий не смешаны, но последние смешаны с эффектами трехфакторных взаимодействий. Все три рассмотренные выше дробные реплики имеют разрешающую способность III.

Кратко остановимся на методе вычисления МНК-оценок параметров функции регрессии вида (1.2) по результатам изме-

рений, проводимых согласно полному факторному плану и его дробным репликам.

Сначала покажем, что при выборе любой функции регрессии вида (1.2) матрица плана F для полного факторного плана 2^m обладает свойством

$$F^T F = N I_s, \quad (1.20)$$

где $N = 2^m$, I_s — единичная матрица порядка s (s — количество параметров функции регрессии).

Столбцы матрицы плана F обозначим через

$$f_i = (f_i(1), \dots, f_i(N))^T, \quad i = 0, \dots, s-1.$$

Поскольку все $f_i(j)$ равны 1 или -1 , то диагональные элементы матрицы $F^T F$ равны $f_i^T f_i = N$ ($i = 0, \dots, s-1$).

Покажем теперь, что все столбцы f_i матрицы F ортогональны. Это эквивалентно тому, что при любом l ($i \leq l \leq m$) и любых $i_1 < \dots < i_l$ выполнено равенство

$$\sum_{j=1}^N x_{i_1}(j) x_{i_2}(j) \dots x_{i_l}(j) = 0, \quad (1.21)$$

где $x_i(j)$ — значение фактора x_i в j -м измерении. Чтобы убедиться в справедливости (1.21), достаточно провести перенумерацию:

$$\begin{aligned} x_{i_1}(j) &= \begin{cases} 1, & j \in \{1, \dots, N/2\}, \\ -1, & j \in \{N/2 + 1, \dots, N\}; \end{cases} \\ x_{i_2}(j) &= \begin{cases} 1, & j \in \{1, \dots, N/4\} \cup \{N/2 + 1, \dots, N - N/4\}, \\ -1, & j \in \{N/4 + 1, \dots, N/2\} \cup \{N - N/4 + 1, \dots, N\}; \end{cases} \\ &\dots \dots \dots \\ x_{i_l}(j) &= \begin{cases} 1, & j \in \{1, \dots, N/2^l\} \cup \dots \cup \{N - N/2^{l-1} + 1, \dots, \\ & \dots, N - N/2^l\}, \\ -1, & j \in \{N/2^l + 1, \dots, N/2^{l-1}\} \cup \dots \cup \{N - N/2^l + \\ & + 1, \dots, N\}. \end{cases} \end{aligned}$$

Тогда левая часть (1.21) будет иметь вид

$$\underbrace{(1 + \dots + 1)}_{N/2^l} - \underbrace{(1 + \dots + 1)}_{N/2^l} + \dots - \underbrace{(1 + \dots + 1)}_{N/2^l}.$$

Таким образом, доказана справедливость (1.20). Отсюда и из вида МНК-оценки ((1.20) из гл. 1) вытекает, что МНК-оценка параметров θ_α имеет вид

$$\hat{\theta}_\alpha = \frac{1}{N} \sum_{j=1}^N f_i(j) y_j. \quad (1.22)$$

где $N = 2^m$; θ_α — один из параметров $\theta_{i_1, i_2, \dots, i_l}$ ($1 \leq l \leq m$) функции регрессии (1.2); I_l — столбец матрицы плана F , соответствующий параметру θ_α ; $y_1, \dots, y_N$ — результаты измерений. Если случайные величины $y_1, \dots, y_N$ некоррелированы и имеют одинаковую дисперсию σ^2 , то из выражения для дисперсионной матрицы МНК-оценок следует, что оценки (1.22) также некоррелированы, а их дисперсии равны σ^2/N .

Дробную реплику 2^{m-p} , примененную для схемы регрессии вида (1.2) с числом параметров, не превосходящим $N = 2^{m-p}$, можно рассматривать как полный факторный эксперимент 2^r ($r = m - p$), примененный к модели того же вида, но с параметрами θ_α , замененными согласно генерирующим соотношениям. Поэтому формулы для расчета МНК-оценок неизвестных параметров θ_α и их дисперсий остаются такими же, как и для полного факторного плана 2^m .

Упражнения.

1. Аналогично рис. 3 изобразите полуреплику 2^{3-1} , определяемую матрицей плана (1.12).

2. Выпишите определяющие соотношения и систему совместных оценок для дробной реплики 2^{3-2} , которая определяется генерирующими соотношениями $x_1 = x_1 x_2 x_3$, $x_2 = x_1 x_3$. Определите разрешающую способность этой дробной реплики.

3. Задайте генерирующие соотношения для дробной реплики 2^{7-3} разрешающей способности: а) III; б) IV.

4. Задайте генерирующие соотношения для полуреплики 2^{7-4} разрешающей способности: а) III; б) IV; в) V; г) VI; д) VII.

§ 2. Латинские планы

2.1. Латинские квадраты. Латинским квадратом порядка n называется расположение символов в виде квадратной таблицы с n строками и n столбцами такое, что каждый символ один раз появляется в каждой строке и один раз — в каждом столбце. Например,

A	B	C	D	E
C	D	E	A	B
E	A	B	C	D
B	C	D	E	A
D	E	A	B	C

Существование латинских квадратов любого порядка следует из примера

A	B	C	D	E
B	C	D	E	A
C	D	E	A	B
D	E	A	B	C
E	A	B	C	D

который обобщается на квадраты любого порядка.

Схема латинского квадрата в статистике возникла на основе сельскохозяйственного эксперимента, типичным примером кото-

рого является следующий. Предположим, что нужно сравнить урожайность пяти сортов пшеницы. Для этого в эксперименте, проводимом по схеме латинского квадрата, прямоугольное поле делят на 25 участков (ячеек), разделенных на пять строк и пять столбцов, причем каждый сорт встречается однажды в каждом столбце и в каждой строке. При этом предполагают, что истинное среднее участка является суммой эффекта строки, эффекта столбца и среднего урожая сорта. Такая схема применима, если среднее участка является суммой среднего урожая и эффекта плодородия, причем имеются колебания плодородия по строкам и столбцам (вызванные, например, различными видами почвы или удобрений).

Число различных латинских квадратов растет с ростом порядка. Новые латинские квадраты могут быть получены из заданного, например, путем перестановки строк, столбцов, а также перестановки символов.

Каждому латинскому квадрату порядка n соответствует латинский план трехфакторного эксперимента, в котором все три фактора имеют n уровней, а общее число измерений равно n^2 . Такое соответствие можно получить, если номер строки отождествить с уровнем первого фактора, номер столбца — с уровнем второго, а символ — с уровнем третьего.

При указанном соответствии латинскому квадрату

1	0
0	1

соответствует изображенная на рис. 3 полуреплика 2^{3-1} полного факторного плана 2^3 , а латинскому квадрату

-1	0	1
0	1	-1
1	-1	0

соответствует план трехфакторного эксперимента, изображенный на рис. 4. На этом рисунке уровни факторов равны $-1, 0, 1$, а все помеченные точки соответствуют полному факторному плану.

2.2. Дисперсионный анализ для латинских квадратов. Предположим, что имеются три фактора a, b, c , принимающих значения на n уровнях, занумерованных числами $1, 2, \dots, n$, и что планом эксперимента является латинский план, причем латинский квадрат, соответствующий плану, выбран из множества всевозможных латинских квадратов порядка n случайным образом. Пусть y_{ijk} — результат измерения при условии, что фактор a находился на уровне i , фактор b — на уровне j , фактор c — на уровне k . Множество из n^2 значений, которые может принимать упорядоченная тройка (i, j, k) , обозначим через Λ .

Преимущество латинского плана по сравнению с полным трехфакторным экспериментом заключается в существенном сокращении числа измерений —

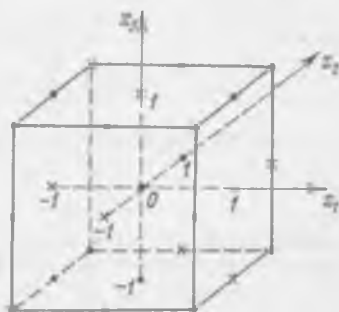


Рис. 4

с n^3 до n^2 . Недостатком является то, что при применении латинских планов числа уровней всех факторов должны быть одинаковыми, а также в том, что дисперсионный анализ для латинских планов проводится при более ограничительных предположениях — необходимо, чтобы взаимодействия между факторами отсутствовали.

Пусть измерения проведены согласно латинскому плану. Предположения о модели, при которых будет проводиться дисперсионный анализ, заключаются в следующем:

$$y_{ijk} = \mu + \alpha_i + \beta_j + \gamma_k + e_{ijk}, \quad (2.1)$$

где $(i, j, k) \in \Lambda$; случайные величины $\{e_{ijk}\}$ независимы и нормально распределены с нулевым средним и конечной неизвестной дисперсией σ^2 ; μ , α_i , β_j , γ_k ($i, j, k = 1, \dots, n$) — неизвестные параметры, называемые *общим средним* и *эффектами факторов* a , b , c соответственно; не ограничивая общности, можно считать, что

$$\sum_{i=1}^n \alpha_i = \sum_{j=1}^n \beta_j = \sum_{k=1}^n \gamma_k = 0. \quad (2.2)$$

поскольку если какая-нибудь из сумм не равна нулю, то ее значение можно включить в μ . Проверяемые гипотезы о независимости результатов измерений от факторов a , b , c записываются в виде

$$\begin{aligned} H_a: \alpha_i &= 0, & i &= 1, \dots, n; \\ H_b: \beta_j &= 0, & j &= 1, \dots, n; \\ H_c: \gamma_k &= 0, & k &= 1, \dots, n \end{aligned}$$

Сначала при рассмотренных предположениях найдем МНК-оценки параметров μ , α_i , β_j , γ_k , минимизируя по указанным переменным функцию

$$Q = \sum_{(i, j, k) \in \Lambda} (y_{ijk} - \mu - \alpha_i - \beta_j - \gamma_k)^2.$$

Приравняв нулю производную $\partial Q / \partial \mu$, получаем

$$\sum_{(i, j, k) \in \Lambda} (y_{ijk} - \mu - \alpha_i - \beta_j - \gamma_k) = 0.$$

Суммируя по всем $(i, j, k) \in \Lambda$, мы суммируем по всем наблюдениям, и для каждого i параметр α_i складывается n раз; следовательно,

$$\sum_{(i, j, k) \in \Lambda} \alpha_i = \sum_i (n\alpha_i) = 0.$$

Аналогично получается и для $\{\beta_j\}$ и $\{\gamma_k\}$. Отсюда получаем, что МНК-оценкой параметра μ является

$$\hat{\mu} = \bar{y} \dots = \frac{1}{n^3} \sum_{(i, j, k) \in \Lambda} y_{ijk}.$$

Приравняв нулю производную $\partial Q / \partial \gamma_k$, получаем

$$\sum_{(i, j) \in \Lambda_k} (y_{ijk} - \mu - \alpha_i - \beta_j - \gamma_k) = 0, \quad (2.3)$$

где Λ_k является множеством таких n пар (i, j) , для которых $(i, j, k) \in \Lambda$. По определению латинского квадрата при каждом k множество Λ_k состоит из таких n пар (i, j) , в которых i и j принимают одно и только одно значение из множества $\{1, \dots, n\}$. Следовательно, выполнены равенства

$$\sum_{(i, j) \in \Lambda_k} \alpha_i = \sum_{i=1}^n \alpha_i = 0, \quad \sum_{(i, j) \in \Lambda_k} \beta_j = \sum_{j=1}^n \beta_j = 0.$$

С учетом этого выражение (2.3) сводится к равенству

$$y_{..k} - \mu - \gamma_k = 0.$$

где

$$y_{..k} = \frac{1}{n} \sum_{(i, j) \in \Lambda_k} y_{ijk}.$$

Таким образом, МНК-оценка параметра γ_k равна

$$\hat{\gamma}_k = y_{..k} - y_{...}$$

Аналогично

$$\hat{\alpha}_i = y_{i..} - y_{...}, \quad \hat{\beta}_j = y_{.j.} - y_{...}$$

где

$$y_{...} = \frac{1}{n^2} \sum_{(i, j, k) \in \Lambda} y_{ijk}$$

а величины $y_{i..}$, $y_{.j.}$ определяются аналогично $y_{..k}$.

Вычисляя теперь остаточную сумму квадратов

$$R_0^2 = \sum_{(i, j, k) \in \Lambda} (y_{ijk} - \hat{\mu} - \hat{\alpha}_i - \hat{\beta}_j - \hat{\gamma}_k)^2,$$

имеем

$$R_0^2 = \sum_{(i, j, k) \in \Lambda} y_{ijk}^2 - \sum_{(i, j, k) \in \Lambda} (\hat{\mu} + \hat{\alpha}_i + \hat{\beta}_j + \hat{\gamma}_k)^2.$$

После элементарных преобразований это равенство переписывается в виде

$$R_0^2 = S_0^2 - S_a^2 - S_b^2 - S_c^2,$$

где

$$S_0^2 = \sum_{(i, j, k) \in \Lambda} y_{ijk}^2 - E, \quad S_a^2 = n \sum_i y_{i..}^2 - E,$$

$$S_b^2 = n \sum_j y_{.j.}^2 - E, \quad S_c^2 = n \sum_k y_{..k}^2 - E,$$

$$E = n^2 y_{...}^2.$$

В модели (2.1) имеется $3n + 1$ параметров, подчиненных трем дополнительным связям (2.2), и, следовательно, $r = 3n - 2$ независимых параметров. Поэтому числом степеней свободы величины R^2 , имеющей χ^2 распределение, является $n^2 - r = n^2 - 3n + 2$.

Пусть теперь наряду с исходными предположениями выполнено

$$H_c: \gamma_k = 0, \quad k = 1, \dots, n.$$

Минимизируя по μ , α_i , β_j функцию

$$Q_c = \sum_{(i, j, k) \in \Lambda} (y_{ijk} - \mu - \alpha_i - \beta_j)^2,$$

получаем, что при выполнении гипотезы H_c МНК-оценки параметров μ , α_i , β_j будут совпадать с полученными выше, а минимум функции Q_c равен

$$\min Q_c = \min Q + S_c^2,$$

где

$$S_c^2 = n \sum_k \gamma_k^2 = n \sum_k y_{..k}^2 - E.$$

Случайная величина S_c^2 имеет χ^2 -распределение с $n - 1$ степенями свободы. Таким образом, при выполнении гипотезы H_c статистика

$$[S_c^2 / (n - 1)] / [R_0^2 / (n^2 - 3n + 2)]$$

имеет F -распределение Фишера с $n - 1$ и $n^2 - 3n + 2$ степенями свободы и может использоваться в качестве тестовой при проверке этой гипотезы.

Аналогично проверяются гипотезы H_a и H_b .

2.3. Ортогональные латинские квадраты. Два латинских квадрата порядка $n \times n$ называются *ортогональными*, если при наложении одного на другой каждая из n^2 пар символов (i, j) ($i, j = 1, \dots, n$) встречается только один раз.

Рассмотрим, например, следующие два латинских квадрата, составленных из латинских и греческих букв:

A	B	C	D	E	α	β	γ	δ	ϵ
C	D	E	A	B	δ	ϵ	α	β	γ
E	A	B	C	D	β	γ	δ	ϵ	α
B	C	D	E	A	ϵ	α	β	γ	δ
D	E	A	B	C	γ	δ	ϵ	α	β

После наложения этих квадратов получаем квадрат

A α	B β	C γ	D δ	E ϵ
C δ	D ϵ	E α	A β	B γ
E β	A γ	B δ	C ϵ	D α
B ϵ	C α	D β	E γ	A δ
D γ	E δ	A ϵ	B α	C β

в котором каждая буква латинского алфавита встречается с каждой буквой греческого лишь один раз.

По аналогии с рассмотренным примером квадраты, получаемые после наложения латинских квадратов, называют *греко-латинскими*.

Греко-латинские квадраты могут быть использованы аналогично латинским для проверки гипотез об отсутствии влияния факторов на результаты измерений в четырехфакторном дисперсионном анализе (конечно, для случая, когда взаимодействия факторов отсутствуют). При использовании таких квадратов достаточно проводить n^2 измерений вместо n^4 , необходимых при проведении стандартного четырехфакторного дисперсионного анализа.

Существование ортогональных латинских квадратов и их построение при данном значении n является вопросом далеко не простым. Так, в 1782 г. Леонард Эйлер предположил, что для всех $n = 4k + 2$ ($k = 1, 2, \dots$) ортогональных латинских квадратов не существует. Предположение Эйлера было опровергнуто только в 1959 г. Оказалось, что ортогональные латинские квадраты существуют для всех n , кроме $n = 2$ и $n = 6$.

Если имеется k попарно ортогональных латинских квадрата (легко показать, что $k \leq n - 1$), то по ним аналогично предыдущему можно проводить $(k + 2)$ -факторный дисперсионный анализ. План эксперимента, построенный по такой системе латинских квадратов, по сравнению с полным многофакторным дисперсионным анализом гораздо более экономичен как с точки зрения количества необходимых измерений, так и с точки зрения вычислений.

2.4. Пример построения полной системы латинских квадратов. Полной системой латинских $n \times n$ -квадратов называется набор из $k = n - 1$ попарно ортогональных латинских $n \times n$ -квадратов.

Между существованием полной системы латинских квадратов и так называемыми полями Галуа существует глубокая связь. Мы рассмотрим эту связь лишь в простейшем случае, когда $n = p$ — простое число.

Пусть $p = 5$. Построим поле вычетов по модулю 5, которое является полем Галуа (о полях Галуа можно подробно прочитать в [43]).

Говорят, что два целых числа l и m сравнимы по модулю 5, если $l - m = 5r$, где r — какое-либо целое число. Это записывается в виде

$$l \equiv m \pmod{5}.$$

Приведенная операция сравнения определяет поле, состоящее из пяти элементов: 0, 1, 2, 3, 4. Составим таблицы сложения и, умножения в этом поле:

Сложение						Умножение					
0	1	2	3	4		1	2	3	4		
1	2	3	4	0		2	4	1	3		
2	3	4	0	1		3	1	4	2		
3	4	0	1	2		4	3	2	1		
4	0	1	2	3							

Рассмотрим латинский квадрат, образованный таблицей сложения, и построим еще три латинских квадрата. Для того чтобы построить l -й латинский квадрат ($l = 2, 3, 4$), строку первого квадрата, начинающуюся с числа m ($m = 0, 1, 2, 3, 4$), заменяем строкой, полученной из первой сложением ее элементов с числом $m \times l$ (имеется в виду сложение и умножение в рассматриваемом поле). Таким образом, получаем еще три латинских квадрата:

$l \times m$						2-й квадрат					
$2 \times 0 = 0$						0	1	2	3	4	
$2 \times 1 = 2$						2	3	4	0	1	
$2 \times 2 = 4$						4	0	1	2	3	
$2 \times 3 = 1$						1	2	3	4	0	
$2 \times 4 = 3$						3	4	0	1	2	
						3-й квадрат					
$3 \times 0 = 0$						0	1	2	3	4	
$3 \times 1 = 3$						3	4	0	1	2	
$3 \times 2 = 1$						1	2	3	4	0	
$3 \times 3 = 4$						4	0	1	2	3	
$3 \times 4 = 2$						2	3	4	0	1	
						4-й квадрат					
$4 \times 0 = 0$						0	1	2	3	4	
$4 \times 1 = 4$						4	0	1	2	3	
$4 \times 2 = 3$						3	4	0	1	2	
$4 \times 3 = 2$						2	3	4	0	1	
$4 \times 4 = 1$						1	2	3	4	0	

Сначала проверим, что полученные квадраты действительно являются латинскими. Предположим, что в i -й строке квадрата, полученного с помощью множителя l ($l=0, 2, 3, 4$), имеется два одинаковых элемента: один в j_1 -м столбце, другой в j_2 -м. В соответствии с правилом образования элементов имеем

$$x_{j_1} + lx_i \equiv x_{j_2} + lx_i \pmod{5},$$

т. е. $x_{j_1} - x_{j_2} = 5r$, где r — целое, а x_{j_1} и x_{j_2} меньше 5. Следовательно, $j_1 = j_2$. Аналогично, если два одинаковых элемента имеются в j -м столбце и в i_1 -й и i_2 -й строках, то

$$x_j + lx_{i_1} \equiv x_j + lx_{i_2} \pmod{5},$$

т. е. $l(x_{i_1} - x_{i_2}) = 5r$. Вследствие того что число 5 простое и $l < 5$, получаем

$$x_{i_1} \equiv x_{i_2} \pmod{5},$$

и, следовательно, $i_1 = i_2$ (так как $x_{i_1} < 5$, $x_{i_2} < 5$).

Покажем теперь, что любые два квадрата, соответствующие множителям l_1 и l_2 ($l_1 > l_2$), ортогональны. Для этого нужно показать, что любая пара элементов при наложении квадратов встречается только один раз. Предположим, что одинаковые пары встретились в клетках (i_1, j_1) и (i_2, j_2) . Тогда

$$x_{j_1} + l_1 x_{i_1} \equiv x_{j_2} + l_1 x_{i_2} \pmod{5},$$

$$x_{j_1} + l_2 x_{i_1} \equiv x_{j_2} + l_2 x_{i_2} \pmod{5},$$

откуда

$$x_{j_1} + l_1 x_{i_1} - x_{j_2} - l_1 x_{i_2} = 5r_1,$$

$$x_{j_1} + l_2 x_{i_1} - x_{j_2} - l_2 x_{i_2} = 5r_2,$$

$$(l_1 - l_2)(x_{i_1} - x_{i_2}) = 5(r_1 - r_2).$$

Поскольку число 5 простое и $l_1 - l_2 < 5$, то $x_{i_1} - x_{i_2}$ делится на 5, и, следовательно, $i_1 = i_2$, так как $x_{i_1} < 5$, $x_{i_2} < 5$. Аналогично получаем $j_1 = j_2$. Таким образом, обе одинаковые пары могут находиться только в одной клетке. Это означает, что квадраты ортогональны.

Точно таким же образом строятся полные системы латинских квадратов для случая, когда число $n = p$ отлично от 5, но является простым.

2.5. Латинские прямоугольники. Использование латинских планов часто бывает неудобным, так как три исследуемых фактора должны иметь одинаковое число уровней. Избежать этого позволяют латинские прямоугольники, которые можно определить как подмножество строк (столбцов) латинского квадрата. Однако не любые указанные подмножества обладают свойствами, делающими удобным соответствующий дисперсионный

анализ. Хорошими (см. § 3) считаются так называемые $n \times k$ -квадраты Юдена ($k < n$), т. е. такие латинские $n \times k$ -прямоугольники, в которых каждая пара символов появляется вместе в $\lambda = k(k-1)/(n-1)$ столбцах. Например, 7×4 -квадратом Юдена является

E	D	G	B	C	A	F
D	C	F	A	B	G	E
A	G	C	E	F	D	B
F	E	A	C	D	B	G

2.6. Латинские кубы. Кроме латинских прямоугольников на практике широко используется еще одно обобщение латинских квадратов — латинские кубы.

Латинский $n \times n \times n$ -куб первого порядка — это такое кубическое размещение n символов, что каждый символ повторяется n^2 раз во всей кубической таблице и ровно n раз в каждой из плоскостей, параллельных трем координатным плоскостям (указанные плоскости будем называть *слоями*). Обычно используют *регулярные латинские кубы*, т. е. такие, у которых все слои являются латинскими квадратами. Например, регулярным латинским $4 \times 4 \times 4$ -кубом первого порядка является куб со слоями:

1	2	3	4
A B C D	B A D C	C D A B	D C B A
B A D C	A B C D	D C B A	C D A B
C D A B	D C B A	A B C D	B A D C
D C B A	C D A B	B A D C	A B C D

Латинский $n \times n \times n$ -куб второго порядка — такое кубическое размещение n^2 элементов, каждый из которых повторяется n раз в кубической таблице и только один раз в каждом слое. Например, латинским $n \times n \times n$ -кубом второго порядка является куб со слоями:

1-й слой	2-й слой	3-й слой
0 4 8	3 7 2	6 1 5
1 5 6	4 8 0	7 2 3
2 3 7	5 6 1	8 0 4

Латинские $n \times n \times n$ -кубы первого порядка могут использоваться аналогично латинским квадратам для проведения четырехфакторного дисперсионного анализа в случае, когда все четыре фактора имеют одинаковое число n уровней, а кубы второго порядка — в случае, если три фактора имеют n уровней, а четвертый имеет n^2 уровней.

Во всех планах, построенных на основе латинских кубов, фигурируют n и n^2 уровней, а число измерений равно n^3 .

Аналогично греко-латинским квадратам могут быть построены греко-латинские кубы, обладающие похожими свойствами.

Разумеется, по аналогии с латинскими кубами можно определить и латинские гиперкубы.

Упражнения.

1. Аналогично рис. 3 изобразите латинский план, соответствующий латинскому квадрату

$$\begin{array}{cc} 0 & 1 \\ 1 & 0 \end{array}$$

2. Аналогично рис. 4 изобразите латинский план, соответствующий латинскому квадрату

$$\begin{array}{ccc} -1 & 1 & 0 \\ 1 & 0 & -1 \\ 0 & -1 & 1 \end{array}$$

3. Постройке греко-латинские квадраты третьего и четвертого порядка.

4. Постройте три попарно ортогональных латинских квадрата седьмого порядка.

5. Укажите, какой из приведенных ниже прямоугольников является квадратом Юдена, а какой не является:

1	2	3	4	5	1	2	3	4	5
3	4	5	1	2	2	3	4	5	1
5	1	2	3	4	3	4	5	1	2
2	3	4	5	1					

6. Является ли латинским кубом куб, все четыре слоя которого представляют собой одинаковые квадраты

1	2	3	4
2	1	4	3
3	4	1	2
4	3	2	1

7. Приведите примеры латинских $2 \times 2 \times 2$ -кубов первого и второго порядка.

§ 3. Неполноблочные планы

3.1. Основные понятия. В § 1, 2 рассматривались три типа факторных планов: полные факторные планы, дробные реплики от них и планы, основанные на комбинаторных конфигурациях типа латинских квадратов. Более общими способами построения факторных планов являются способы, основанные на построении так называемых блок-схем.

Двумерные блок-схемы определяются следующим образом.

Пусть заданы два конечных множества

$$A = \{a_1, \dots, a_v\}, \quad B = \{B_1, \dots, B_b\}.$$

Говорят, что эти множества порождают *инцидентную структуру*, если между их элементами установлено отношение инцидентности (принадлежности) $a_i \in B_j$. *Двумерной блок-схемой* называется инцидентная структура, порождаемая множествами A, B ; при этом a_i ($i = 1, \dots, v$) принято называть *элементами, образцами* или *способами обработки*, а B_j ($j = 1, \dots, b$) — *блоками*.

Понятие двумерной блок-схемы может быть обобщено на многомерный случай. При этом в определении возрастает число

множеств и отношений инцидентности. Трехмерной блок-схемой является, например, латинский квадрат размера $v \times v$, в котором рассматриваются отношения инцидентности множества элементов $A = \{a_1, \dots, a_v\}$ к столбцам B_j (множество B) и строкам (множество C). Греко-латинский квадрат является примером четырехмерной блок-схемы, в которой дополнительно к указанным двум отношениям инцидентности добавляется еще одно — отношение множества латинских букв к множеству греческих. Четырехмерную блок-схему индуцирует и латинский куб.

Ниже будут рассматриваться только двумерные блок-схемы, причем слово «двумерный» будет опускаться.

Параметрами произвольной блок-схемы являются: v — число элементов; b — число блоков; k_j ($j = 1, \dots, b$) — число элементов в блоке B_j ; r_i ($i = 1, \dots, v$) — число блоков, содержащих элемент a_i ; $\lambda_{i_1 i_2}$ ($i_1, i_2 = 1, \dots, v$) — число блоков, содержащих пару элементов $\{a_{i_1}, a_{i_2}\}$.

Задание блок-схемы эквивалентно заданию блочного плана двухфакторного эксперимента, в котором элементы a_i соответствуют уровням первого фактора, а блоки B_j — уровням второго. Следовательно, каждое утверждение о блок-схеме можно истолковывать как утверждение о блочном плане двухфакторного эксперимента.

Блок-схема называется *правильной*, если все блоки имеют один и тот же размер k , т. е. $k_j = k$ ($j = 1, \dots, b$).

Правильная блок-схема называется *полной*, если $k = v$, и *неполной*, если $k < v$. Основным интерес представляют неполные блок-схемы, поскольку полные соответствуют полному факторному эксперименту. Выбор неполной блок-схемы может быть связан как с соображениями экономии числа измерений, так и с невозможностью выбора $k = v$; например, при сравнении v ($v > 4$) марок различных автомобильных шин естественный блок состоит из четырех колес автомобиля.

Блок-схема называется *равноповторной*, если $r_i = r$ ($i = 1, \dots, v$). Правильная равноповторная блок-схема называется *симметричной*, если $v = b$. Правильная равноповторная схема называется *сбалансированной*, если любая пара элементов принадлежит одному и тому же числу λ блоков, т. е. если выполнено условие

$$\lambda_{i_1 i_2} = \lambda, \quad i_1, i_2 = 1, 2, \dots, v. \quad (3.1)$$

Правильная равноповторная схема называется *частично сбалансированной* при выполнении некоторого ослабления (3.1).

Существуют разные способы задания блочного плана и соответствующей ему блок-схемы. Приведем наиболее известные. Способ 1 — задание в виде набора блоков $B = \{B_1, \dots, B_b\}$. Способ 2 — задание в виде таблицы с b столбцами, где столбец соответствует блоку. Если блок-схема является правильной, то

получающаяся таблица будет прямоугольной (имеет k строк и b столбцов). Этот способ наиболее экономный. Способ 3 — задание в виде матрицы инцидентий $N = \|n_{ij}\|$ ($i = 1, \dots, v$; $j = 1, \dots, b$), где

$$n_{ij} = \begin{cases} 1, & \text{если } a_i \in B_j, \\ 0, & \text{если } a_i \notin B_j. \end{cases}$$

Способ 4 — задание в виде таблицы факторного эксперимента, по сути аналогичный предыдущему. Способ 5 — задание в виде матрицы плана F вида (1.2) из гл. 1. Этот способ задания блочного плана требует предварительного задания вида модели, в которой план используется.

Стандартной моделью, в которой применяются блочные планы, является модель линейного регрессионного анализа, согласно которой результат измерения i -го элемента ($i = 1, \dots, v$) в j -м блоке ($j = 1, \dots, b$) представляет собой случайную величину

$$y_{ij} = \mu + \alpha_i + \beta_j + \varepsilon_{ij}, \quad (3.2)$$

где μ , α_i , β_j — неизвестные параметры; ε_{ij} — случайные ошибки, предполагающиеся некоррелированными и имеющими нулевое среднее и дисперсию σ^2 . Параметр μ в (3.2) называется *общим средним*, α_i — *эффектом i -го элемента (образца)*, β_j — *эффектом j -го блока (столбца)*. Число измерений случайной величины y_{ij} равно

$$N = \sum_{i=1}^v k_i = \sum_{j=1}^b r_j.$$

Параметры модели (3.2) не являются независимыми, они удовлетворяют соотношениям

$$\sum_{i=1}^v \alpha_i = \sum_{j=1}^b \beta_j = 0.$$

Схему измерений (3.2) можно записать в виде классической линейной регрессионной модели $(Y, F\theta, \sigma^2 I_N)$, где

$$\theta = (\mu, \alpha_1, \dots, \alpha_v, \beta_1, \dots, \beta_b)^T$$

есть вектор неизвестных параметров, а матрица F имеет N строк, в каждой из которых $v + b - 2$ нулей и три единицы, причем единицами являются первая, $(i + 1)$ -я и $(v + 1 + j)$ -я компоненты, где i и j — соответственно номера элемента и блока, для которых проводилось измерение (3.2).

Некоторые блочные планы могут быть заданы также в виде набора реплик, где каждая реплика состоит из нескольких блоков и содержит одинаковое число раз все элементы a_i ($i = 1, \dots, v$). Задание блочных планов в виде реплик встречается, когда проводимые измерения неоднородны во времени. В этих случаях общее время, отведенное на эксперимент, делится на несколько равных промежутков, в каждом из которых

элементы измеряются одинаковое число раз. Тем самым элементы равномерно распределяются во времени.

Пример 3.1. Блочный план с параметрами $v=b=4$, $k=r=3$, $\lambda=2$ можно представить следующими способами.

Способ 1 (набор блоков):

$$B_1 = \{1, 2, 3\}, \quad B_2 = \{2, 3, 4\}, \quad B_3 = \{3, 4, 1\}, \quad B_4 = \{4, 1, 2\}.$$

Способ 2 (прямоугольная таблица, в которой столбцы соответствуют блокам):

$$\begin{array}{cccc} 1 & 2 & 3 & 4 \\ 2 & 3 & 4 & 1 \\ 3 & 4 & 1 & 2 \end{array}$$

Способ 3 (матрица инциденций):

$$N = \begin{bmatrix} 1 & 0 & 1 & 1 \\ 1 & 1 & 0 & 1 \\ 1 & 1 & 1 & 0 \\ 0 & 1 & 1 & 1 \end{bmatrix}.$$

Способ 4 (таблица факторного эксперимента):

Элементы	Блоки			
	B_1	B_2	B_3	B_4
1	+		+	+
2	+	+		+
3	+	+	+	
4		+	+	+

Способ 5 (матрица плана F для модели (3.2)). Пусть измерения (3.2) проводятся в том порядке, в каком они выписаны в способе 1. Тогда матрица плана F модели $(Y, F\theta, \sigma^2 I_N)$ имеет вид

$$F = \begin{bmatrix} 1 & 1 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 1 & 0 & 1 & 0 & 0 & 1 & 0 & 0 & 0 \\ 1 & 0 & 0 & 1 & 0 & 1 & 0 & 0 & 0 \\ 1 & 0 & 1 & 0 & 0 & 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 1 & 0 & 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 & 1 & 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 1 & 0 & 0 & 0 & 1 & 0 \\ 1 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 1 \\ 1 & 1 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 1 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 1 \\ 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \\ 1 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix}.$$

Еще один способ представлен на рис. 5.

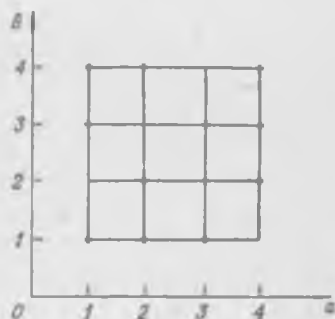


Рис. 5

Пример 3.2. Блочный план с параметрами $v=9$, $b=12$, $k=3$, $r=4$, $\lambda=1$ представим следующими способами.

Способ 1 (блоки):

$$\begin{aligned} B_1 &= \{1, 2, 3\}, & B_2 &= \{4, 5, 6\}, & B_3 &= \{7, 8, 9\}, \\ B_4 &= \{1, 4, 7\}, & B_5 &= \{2, 5, 8\}, & B_6 &= \{3, 6, 9\}, \\ B_7 &= \{1, 5, 9\}, & B_8 &= \{2, 6, 7\}, & B_9 &= \{3, 4, 8\}, \\ B_{10} &= \{1, 6, 8\}, & B_{11} &= \{2, 4, 9\}, & B_{12} &= \{3, 5, 7\}. \end{aligned}$$

Способ 2 (таблица):

1	4	7	1	2	3	1	2	3	1	2	3
2	5	8	4	5	6	5	6	4	6	4	5
3	6	9	7	8	9	9	7	8	8	9	7

Способ 3 (матрица инцидентий):

$$N = \begin{bmatrix} 1 & 0 & 0 & 1 & 0 & 0 & 1 & 0 & 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 & 1 & 0 & 0 & 1 & 0 & 0 & 1 & 0 \\ 1 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 1 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 1 \\ 0 & 1 & 0 & 0 & 0 & 1 & 0 & 1 & 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 1 & 1 & 0 & 0 & 0 & 1 & 0 \end{bmatrix}.$$

Способ 4 (таблица факторного эксперимента):

Элементы	Блоки											
	B_1	B_2	B_3	B_4	B_5	B_6	B_7	B_8	B_9	B_{10}	B_{11}	B_{12}
1	+			+			+			+	+	
2	+				+			+				
3	+					+			+			+
4		+		+					+		+	
5		+			+		+					+
6		+				+				+		+
7			+	+		+		+				+
8			+		+				+	+		
9			+			+	+				+	

Способ 5 приводить не будем ввиду громоздкости матрицы плана (ее порядок — 36×22). Рассматриваемый план представим на рис. 6 и в виде набора реплик:

$$\begin{array}{llll} \text{Реплика I} & \text{Реплика II} & \text{Реплика III} & \text{Реплика IV} \\ B_1 = \{1, 2, 3\} & B_4 = \{1, 4, 7\} & B_7 = \{1, 5, 9\} & B_{10} = \{1, 6, 8\} \\ B_2 = \{4, 5, 6\} & B_5 = \{2, 5, 8\} & B_8 = \{2, 6, 7\} & B_{11} = \{2, 4, 9\} \\ B_3 = \{7, 8, 9\} & B_6 = \{3, 6, 9\} & B_9 = \{3, 4, 8\} & B_{12} = \{3, 5, 7\} \end{array}$$

Блочный план, соответствующий неполной сбалансированной блок-схеме, называется ВВВ-планом (от английского термина balanced incomplete block design).

Блочный план, соответствующий неполной частично сбалансированной блок-схеме, называется РВІВ-планом (от partially balanced incomplete block design).

Планы, приведенные в примерах 3.1 и 3.2, являются ВІВ-планами, а план из примера 3.1 еще и *симметричен* (т. е. ему соответствует симметричная блок-схема).

ВІВ-планы часто применяются в практическом эксперименте, поскольку обладают рядом свойств оптимальности, простотой проведения дисперсионного анализа и тем, что при оценивании по МНК эффектов элементов $\alpha_1, \dots, \alpha_v$ дисперсии всех оценок и все ковариации этих оценок равны между собой. Как показано в п. 3.2, ВІВ-планы могут быть построены далеко не при любых значениях параметров v, b, k, r, λ . В случаях, когда такое построение невозможно, используются другие планы, например РВІВ-планы, для которых процедура дисперсионного анализа также не слишком сложна, а дисперсии МНК-оценок эффектов элементов тоже равны.

Мы не будем останавливаться на статистических процедурах, связанных с указанными планами, а в п. 3.3 опишем простой, но типичный способ построения ВІВ-планов (существует несколько десятков методов построения ВІВ-планов и еще большее число методов построения РВІВ-планов).

3.2. Некоторые свойства параметров неполных сбалансированных блок-схем. Сбалансированная неполная блок-схема (ВІВ-схема) характеризуется пятью параметрами — числом элементов v , числом блоков b , числом повторений каждого элемента в схеме r , размером блоков k и числом λ повторений каждой пары элементов в блоках.

Выведем некоторые соотношения, которым должны удовлетворять параметры v, b, r, k, λ произвольной ВІВ-схемы.

Во-первых, поскольку блок-схема является неполной, то

$$k < v. \quad (3.3)$$

Во-вторых,

$$bk = vr, \quad (3.4)$$

поскольку общее число N инцидентий в блок-схеме, соответствующее общему числу измерений в факторном эксперименте, может быть представлено двумя способами: с одной стороны,

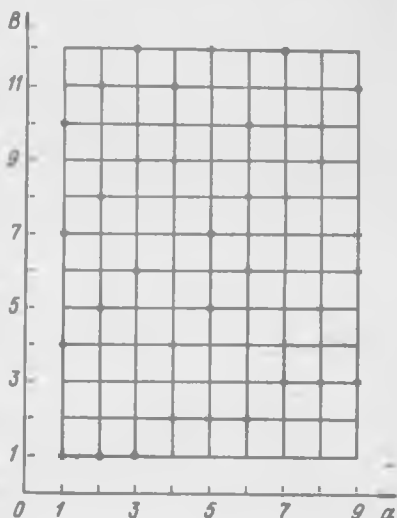


Рис. 6

$N = bk$ (каждый из b блоков содержит по k элементов) и, с другой стороны, $N = vr$ (каждый из v элементов принадлежит ровно r блокам).

Сосчитаем теперь двумя способами число пар элементов, в которые входит фиксированный элемент a_1 . В каждом блоке, в котором содержится a_1 , содержится еще $k - 1$ элементов, и, следовательно, этот элемент входит в $r(k - 1)$ пар. С другой стороны, элемент a_1 входит в пары со всеми остальными $v - 1$ элементами, причем каждая пара встречается λ раз; следовательно, элемент a_1 встречается в $\lambda(v - 1)$ парах. Доказано, таким образом, что

$$r(k - 1) = \lambda(v - 1). \quad (3.5)$$

Далее, если $N = \|n_{ij}\|$ ($i = 1, \dots, v; j = 1, \dots, b$) — матрица инцидентий ВВВ-схемы, то

$$NN^T = \begin{bmatrix} r & \lambda & \dots & \lambda \\ \lambda & r & \dots & \lambda \\ \vdots & \vdots & \ddots & \vdots \\ \lambda & \lambda & \dots & r \end{bmatrix} = (r - \lambda)I_v + \lambda J_v, \quad (3.6)$$

где I_v — единичная $v \times v$ -матрица, а J_v является $v \times v$ -матрицей, состоящей из единиц. Соотношение (3.6) эквивалентно требованию о том, что каждый элемент встречается в схеме r раз, а каждая пара элементов — λ раз.

Лемма 3.1. *Определитель матрицы (3.6) равен*

$$\det NN^T = (r - \lambda)^{v-1} (v\lambda - \lambda + r). \quad (3.7)$$

Доказательство. Вычтем первый столбец матрицы (3.6) из всех остальных, а затем прибавим вторую, третью, ..., v -ю строки к первой. Тогда все элементы, расположенные выше главной диагонали получившейся матрицы, равны нулю; на главной диагонали первый элемент есть $r + (v - 1)\lambda$, а остальные равны $r - \lambda$. Отсюда и следует (3.7). Лемма доказана.

Если $r = \lambda$, то блок-схема является полной, поскольку при $r = \lambda$ каждый элемент a_i может появляться в блоке только вместе с каждым другим элементом a_j и, следовательно, каждый блок содержит все v элементов. Из (3.7) теперь следует, что для ВВВ-схемы

$$\text{rang } NN^T = v. \quad (3.8)$$

С другой стороны, ранг $v \times b$ -матрицы N не превосходит b . Поскольку ранг произведения матриц не превосходит ранга сомножителя (см. теорему 1.13 из приложения 1), то

$$\text{rang } NN^T \leq \text{rang } N \leq b.$$

Отсюда и из (3.8) получаем неравенство

$$v \leq b, \quad (3.9)$$

которое называется *неравенством Фишера*. Следствием (3.4) и (3.9) является неравенство

$$k \leq r. \quad (3.10)$$

Выведем еще одно ограничение, которому должны удовлетворять параметры ВВВ-схем.

Теорема 3.1. *Если числа v и b равны и являются четными, то $k - \lambda$ есть точный квадрат.*

Доказательство. Так как $b = v$, то N — квадратная матрица, и из формулы (3.7) получаем

$$(\det N)^2 = (k - \lambda)^{v-1} (v\lambda - \lambda + k). \quad (3.11)$$

Равенство (3.5) с учетом соотношения $k = r$, которое следует из (3.4) и $b = v$, переписывается в виде

$$k(k - 1) = \lambda(v - 1).$$

Отсюда

$$v\lambda - \lambda + k = k^2,$$

т. е. второй множитель в правой части (3.11) является точным квадратом. Левая часть (3.11) — точный квадрат в силу того, что матрица N состоит из целых чисел и $\det N$ — целое число. Следовательно, первый множитель в правой части (3.11), т. е. $(k - \lambda)^{v-1}$, также является точным квадратом. Утверждение теоремы вытекает из четности v . Теорема доказана.

Из доказанной теоремы следует, в частности, что ВВВ-схемы существуют не для всех значений параметров v, b, r, k, λ , удовлетворяющих равенствам (3.4), (3.5) и неравенствам (3.3), (3.9), (3.10). Так, указанным соотношениям удовлетворяют числа $v = b = 22, r = k = 7, \lambda = 2$. Но так как v четно, а $k - \lambda = 5$ не является точным квадратом, ВВВ-схемы с указанными значениями параметров не существует.

3.3. Построение неполных сбалансированных схем с помощью проективных геометрий над полями Галуа. В данном пункте рассматривается один из относительно простых способов построения неполных сбалансированных схем (ВВВ-схем), определяемых параметрами v, b, k, r, λ .

Для построения ВВВ-схем используем конструкцию, основанную на связи этих схем с проективными геометриями над полями Галуа. Будем рассматривать лишь простейший случай, когда поле Галуа содержит p чисел (p — простое число). В этом случае, как указано в п. 2.4, полем Галуа является поле вычетов по модулю p , элементами этого поля являются числа $0, 1, \dots, p - 1$, а под сложением и умножением понимается сложение и умножение по модулю p . Указанное поле Галуа обозначим через $GF(p)$, а определяемую ниже m -мерную проективную геометрию над этим полем Галуа — через $PG(m, p)$.

В пространстве $PG(m, p)$ точкой является совокупность $m+1$ координат $g_1, \dots, g_{m+1}$, принадлежащих полю Галуа $GF(p)$ и не равных одновременно нулю. Две точки $g = (g_1, \dots, g_{m+1})^T$ и $g' = (g'_1, \dots, g'_{m+1})^T$ считаются совпадающими, если найдется такое $\lambda \in GF(p)$, $\lambda \neq 0$, что для всех $i = 1, \dots, m+1$ выполнено равенство $g_i = \lambda g'_i$ (т. е. координаты точек пропорциональны). Точка, все координаты которой равны нулю, не входит в пространство $PG(m, p)$.

Прямой в пространстве $PG(m, p)$ является множество точек g вида

$$g = \lambda_1 g^{(1)} + \lambda_2 g^{(2)},$$

где $g^{(1)}$ и $g^{(2)}$ — различные точки из $PG(m, p)$, $\lambda_1, \lambda_2 \in GF(p)$. Прямая, следовательно, образована множеством точек с координатами

$$\lambda_1 g^{(1)}_1 + \lambda_2 g^{(2)}_1, \dots, \lambda_1 g^{(1)}_{m+1} + \lambda_2 g^{(2)}_{m+1},$$

где $g^{(j)}_i$ ($i = 1, \dots, m+1$; $j = 1, 2$) — координаты точки $g^{(j)}$.

По аналогии с прямой q -мерная плоскость определяется как множество точек g вида

$$g = \lambda_1 g^{(1)} + \lambda_2 g^{(2)} + \dots + \lambda_{q+1} g^{(q+1)},$$

где $\lambda_i \in GF(p)$ ($i = 1, \dots, q+1$), $g^{(1)}, \dots, g^{(q+1)}$ — линейно независимые точки из $PG(m, p)$.

Установим несколько простых свойств проективных геометрий над полями Галуа.

Из соображений симметрии получаем, что в $PG(m, p)$ через каждую точку проходит одинаковое число v q -мерных плоскостей (точки мы будем ассоциировать с элементами блок-схемы, а q -мерные плоскости — с блоками). Из соображений симметрии получаем также, что все q -мерные плоскости содержат одинаковое число точек k (т. е. каждый элемент встречается в блоках одинаковое число k раз). Наконец, из тех же соображений через каждую прямую в $PG(m, p)$ проходит одинаковое число λ q -мерных плоскостей (это соответствует тому, что любая пара элементов встречается в блоках одинаковое число λ раз). Найдем выражения для параметров, получающихся в сформулированной схеме.

Подсчитаем число точек пространства $PG(m, p)$. Каждая координата может принимать p различных значений, а каждая точка имеет $m+1$ координат. Исключая точку со всеми нулевыми координатами, имеем $p^{m+1} - 1$ точек. Теперь учтем то, что некоторые из этих $p^{m+1} - 1$ точек совпадают — точка с координатами λg_i идентична точке с координатами g_i . Поскольку $\lambda \in GF(p)$ и $\lambda \neq 0$, то можно взять $p-1$ различных значений λ , откуда следует, что одна точка представляется $p-1$ различ-

ными формулами. Поэтому число различных точек пространства $PG(m, p)$ равно

$$v = v(m, p) = \frac{p^{m+1} - 1}{p - 1} = 1 + p + \dots + p^m. \quad (3.12)$$

Для того чтобы подсчитать, сколько точек принадлежит q -мерной плоскости в $PG(m, p)$ (в частности, при $q = 1$ — прямой), докажем следующее вспомогательное утверждение.

Лемма 3.2. q -мерная плоскость ($q < m$) проективной геометрии $PG(m, p)$ является проективной геометрией $PG(q, p)$.

Доказательство. Согласно определению, q -мерная плоскость определяется с помощью $q + 1$ линейно независимых точек из $PG(m, p)$, т. е. таких точек $g^{(1)}, \dots, g^{(q+1)}$, что если

$$\lambda_1 g^{(1)} + \dots + \lambda_{q+1} g^{(q+1)} = 0$$

и $\lambda_1, \dots, \lambda_{q+1} \in GF(p)$, то $\lambda_1 = \lambda_2 = \dots = \lambda_{q+1} = 0$. При этом q -мерная плоскость состоит из точек вида

$$\mu_1 g^{(1)} + \dots + \mu_{q+1} g^{(q+1)}, \quad (3.13)$$

где $\mu_1, \dots, \mu_{q+1} \in GF(p)$ и не все μ_i равны нулю одновременно. Будем считать координатами точек вида (3.13) величины $\mu_1, \dots, \mu_{q+1}$; при этом две точки с координатами $\mu = (\mu_1, \dots, \mu_{q+1})^T$ и $\mu' = (\mu'_1, \dots, \mu'_{q+1})^T$ должны совпадать тогда и только тогда, когда существует такое $v \in GF(p)$, $v \neq 0$, что $\mu_i = v\mu'_i$ для всех $i = 1, \dots, q + 1$.

Из свойств $PG(p, m)$ сразу получаем, что если это соотношение выполняется, то точки с координатами μ и μ' совпадают. С другой стороны, если точки с координатами μ и μ' совпадают, то при любом $i = 1, \dots, m + 1$ для всех $q + 1$ координат этих точек в исходном пространстве $PG(p, m)$ выполнено равенство

$$\begin{aligned} \mu_1 g_i^{(1)} + \mu_2 g_i^{(2)} + \dots + \mu_{q+1} g_i^{(q+1)} = \\ = v(\mu'_1 g_i^{(1)} + \mu'_2 g_i^{(2)} + \dots + \mu'_{q+1} g_i^{(q+1)}). \end{aligned}$$

Отсюда следует, что у точки

$$(\mu_1 - v\mu'_1) g^{(1)} + (\mu_2 - v\mu'_2) g^{(2)} + \dots + (\mu_{q+1} - v\mu'_{q+1}) g^{(q+1)}$$

все координаты нулевые. Используя линейную независимость точек $g^{(1)}, \dots, g^{(q+1)}$, имеем

$$\mu_1 - v\mu'_1 = \mu_2 - v\mu'_2 = \dots = \mu_{q+1} - v\mu'_{q+1} = 0.$$

Это эквивалентно требуемому соотношению $\mu = v\mu'$. Лемма доказана.

В качестве следствия доказанной леммы и формулы (3.12) получаем выражение для числа точек q -мерной плоскости:

$$k = k(p, q) = 1 + p + p^2 + \dots + p^q. \quad (3.14)$$

Теперь рассмотрим способ построения q -мерных плоскостей и определим, сколько существует таких плоскостей.

Сначала произвольным образом выберем первую точку $g^{(1)}$. В силу (3.12) количество способов такого выбора равно $1 + p + \dots + p^m$. Вторая точка $g^{(2)}$ должна быть отлична от первой и поэтому может быть выбрана $p + p^2 + \dots + p^m$ способами. Третья точка $g^{(3)}$ должна быть линейно независимой от первых двух, т. е. отличной от точек прямой

$$\lambda_1 g^{(1)} + \lambda_2 g^{(2)},$$

которая содержит в силу (3.14) $1 + p$ точек. Следовательно, число способов выбора точки $g^{(3)}$ равно

$$p^2 + p^3 + \dots + p^m.$$

Аналогично j -ю точку $g^{(j)}$ ($j \leq q + 1$) можно выбирать из $(1 + p + \dots + p^m) - (1 + p + \dots + p^{j-2}) = p^{j-1} + p^j + \dots + p^m$ точек пространства $PG(m, p)$, которые не принадлежат $(j-2)$ -мерной плоскости, проходящей через точки $g^{(1)}, \dots, g^{(j-1)}$. Таким образом, получаем

$$(1 + p + \dots + p^m)(p + \dots + p^m) \dots (p^q + \dots + p^m) \quad (3.15)$$

q -мерных плоскостей, определяемых по формуле (3.13). Эти плоскости, однако, не все различны, поскольку одна и та же плоскость может быть получена разными способами.

Действительно, каждая q -мерная плоскость содержит

$$1 + p + \dots + p^q$$

точек, а определяется эта плоскость через $q + 1$ точек, выбранных из них. Следовательно, эту плоскость мы получим столько раз, сколько можно найти наборов $q + 1$ линейно независимых точек в q -мерной плоскости (порядок точек в наборах не имеет значения). В силу леммы 3.2 количество таких наборов определяется по формуле (3.15) с заменой m на q ; поэтому общее число различных q -мерных плоскостей в $PG(m, p)$ равно

$$b = b(m, q, p) = \frac{(1 + p + \dots + p^m)(p + \dots + p^m) \dots (p^q + \dots + p^m)}{(1 + p + \dots + p^q)(p + \dots + p^q) \dots p^q}. \quad (3.16)$$

Итак, поставив в соответствие точкам проективной геометрии $PG(m, p)$ элементы двумерной блок-схемы, а q -мерным плоскостям этой геометрии — блоки этой схемы, мы показали, что схема будет сбалансированной, и нашли выражения для трех ее характеристик — v , k , b (формулы (3.12), (3.14), (3.16)). Две другие характеристики — r и λ — находятся по формулам (3.4), (3.5).

Пример 3.3. Построение ВІВ-схемы с помощью двумерных плоскостей проективной геометрии $PG(3, 2)$. В приведенных

обозначениях $m = 3$, $p = 2$, $q = 2$. По формулам (3.12), (3.14), (3.16), (3.4), (3.5) находим все характеристики ВІВ-схемы:

$$v = 1 + 2 + 2^2 + 2^3 = 15,$$

$$k = 1 + 2 + 2^2 = 7,$$

$$b = b(3, 2, 2) = 15,$$

$$r = 7 \cdot 15 / 15 = 7,$$

$$\lambda = 7 \cdot 6 / 14 = 3.$$

Точки пространства $PG(3, 2)$, соответствующие элементам блок-схемы, будем обозначать числами. Эти точки определяются своими четырьмя координатами, соответствующими двоичному разложению чисел:

$$\begin{aligned} 8 &= (1, 0, 0, 0)^T, & 10 &= (1, 0, 1, 0)^T, & 14 &= (1, 1, 1, 0)^T, \\ 4 &= (0, 1, 0, 0)^T, & 9 &= (1, 0, 0, 1)^T, & 13 &= (1, 1, 0, 1)^T, \\ 2 &= (0, 0, 1, 0)^T, & 6 &= (0, 1, 1, 0)^T, & 11 &= (1, 0, 1, 1)^T, \\ 1 &= (0, 0, 0, 1)^T, & 5 &= (0, 1, 0, 1)^T, & 7 &= (0, 1, 1, 1)^T, \\ 12 &= (1, 1, 0, 0)^T, & 3 &= (0, 0, 0, 1)^T, & 15 &= (1, 1, 1, 1)^T. \end{aligned}$$

Построим теперь по указанной выше схеме двумерные плоскости (соответствующие блокам ВІВ-схемы), проводя эти плоскости через две точки, выбираемые каждый раз минимально возможными (при интерпретации их как чисел), и отыскивая остальные точки плоскостей; при этом помним, что каждое число должно встречаться семь раз, а пары чисел — по три раза. Общее количество плоскостей (т. е. блоков) равно 15, каждой плоскости соответствует в приведенной ниже таблице столбец:

1	1	1	1	1	1	1	2	2	2	2	3	3	3	3
2	2	2	4	4	6	6	4	4	5	5	4	4	5	5
3	3	3	5	5	7	7	6	6	7	7	7	7	6	6
4	8	12	8	10	8	10	8	9	8	9	8	9	8	9
5	9	13	9	11	9	11	10	11	10	11	11	10	11	10
6	10	14	12	14	14	12	12	13	13	12	12	13	13	12
7	11	15	13	15	15	13	14	15	15	14	15	14	14	15

Этой таблицей и задается искомая ВІВ-схема.

Упражнения.

1. Укажите, какие из приведенных блок-схем являются сбалансированными и какие частично сбалансированными:

$$\begin{aligned} \text{а) } & \begin{matrix} 0 & 1 & 2 & 3 & 4 & 5 & 6 \\ 1 & 2 & 3 & 4 & 5 & 6 & 0 \\ 3 & 4 & 5 & 6 & 0 & 1 & 2 \end{matrix} & \text{б) } & \begin{matrix} 1 & 1 & 1 & 1 & 4 & 3 & 2 & 2 \\ 2 & 4 & 3 & 2 & 5 & 4 & 3 & 5 \\ 5 & 5 & 4 & 3 & 6 & 6 & 6 & 6 \end{matrix} \\ \text{в) } & \begin{matrix} 1 & 4 & 7 & 1 & 2 & 3 & 1 & 2 & 3 \\ 2 & 5 & 8 & 4 & 5 & 6 & 5 & 6 & 4 \\ 3 & 6 & 9 & 7 & 8 & 9 & 9 & 7 & 8 \end{matrix} \end{aligned}$$

2. Выпишите матрицы инцидентий для блок-схем из предыдущего упражнения

3. Постройте сбалансированные блок-схемы с параметрами $v = 4$, $k = 2$, $r = 3$, $b = 6$, $\lambda = 1$ и $v = 5$, $k = 2$, $r = 4$, $b = 10$, $\lambda = 1$.

4. Покажите, что система нормальных уравнений при оценивании параметров $\alpha = (\alpha_1, \dots, \alpha_v)^T$ схемы регрессии (3.2) в случае использования ВІВ-плана с матрицей инцидентий N имеет вид $C\alpha = Q$, где

$$C = rI_v - k^{-1}NN^T, \quad Q = Y_A - k^{-1}NY_B.$$

Y_A есть v -вектор, компонентами которого являются суммы результатов измерений y_{ij} элементов a_i ($i = 1, \dots, v$), Y_B есть b -вектор, компонентами которого являются суммы результатов измерений по блокам.

5. Постройте ВІВ-план с помощью прямых проективной геометрии $PG(3, 2)$.

6. Чему равно число двумерных плоскостей проективной геометрии $PG(7, 3)$?

7. Докажите, что всевозможные комбинации из v элементов по k образуют блоки ВІВ-схемы с параметрами v , k , $b = C_v^k$, $r = C_v^{k-1}$, $\lambda = C_v^{k-2}$ ($v \geq 3$, $k \geq 2$).

Глава 3

КЛАССИЧЕСКАЯ ТЕОРИЯ ПЛАНИРОВАНИЯ ЭКСПЕРИМЕНТА ПО ОЦЕНИВАНИЮ ПАРАМЕТРОВ ЛИНЕЙНЫХ РЕГРЕССИОННЫХ МОДЕЛЕЙ

В данной главе изложена классическая теория планирования эксперимента по оцениванию параметров линейных регрессионных моделей, называемая также теорией планирования регрессионного эксперимента. Начала этой теории были заложены в работах видного американского математика Дж. Кифера, который, в частности, ввел в рассмотрение понятие непрерывного плана эксперимента.

В § 1 вводятся фундаментальные понятия. В § 2 рассматриваются теоремы эквивалентности. В § 3 на ряде примеров показано, что с помощью теорем эквивалентности удастся изучить свойства оптимальных планов, а в ряде случаев и аналитически построить их. Теоремы эквивалентности являются также основой описанных в § 4 алгоритмов приближенного построения оптимальных планов, суть которых состоит в редукции размерности экстремальной задачи (исходная задача оптимизации высокой размерности сводится к последовательности задач оптимизации меньшей размерности).

§ 1. Основные понятия и вспомогательные результаты

1.1. План эксперимента. Пусть X — некоторое компактное подмножество конечномерного евклидова пространства ^{*}) R^k . Это подмножество будем называть *множеством планирования*. Точки множества X иногда называют *условиями проведения измерений*. В данной главе будет рассматриваться схема проведения измерений, согласно которой результат измерения $y(x_i)$ в точке $x_i \in X$ представляется в виде

$$y_i = y(x_i) = \sum_{l=1}^m \theta_l f_l(x_i) + \varepsilon_i, \quad (1.1)$$

^{*}) На самом деле для справедливости большинства результатов достаточно предполагать, что X — ограниченное замкнутое подмножество некоторого метрического пространства.

где $\theta_1, \dots, \theta_m$ — неизвестные параметры; $f_1(x), \dots, f_m(x)$ — известные измеримые функции на X , которые будем называть *базисными функциями* (в дальнейшем будем предполагать их непрерывность); $e_j = e(x_j)$ — случайные ошибки, удовлетворяющие условиям: а) центрированности ($E e_j = 0$), б) некоррелированности ($E e_i e_j = 0$ при различных измерениях в точках $x_i, x_j \in X$, где $i \neq j$, но условия проведения измерений x_i и x_j могут совпадать), в) равноточности ($E e_j^2 = \sigma^2 < \infty$ для любой точки $x_j \in X$).

В схеме измерений (1.1) функцией регрессии является

$$\eta(x) = E y(x) = \sum_{i=1}^m \theta_i f_i(x) = \theta^T f(x), \quad (1.2)$$

где $\theta = (\theta_1, \dots, \theta_m)^T$ — вектор неизвестных параметров, $f(x) = (f_1(x), \dots, f_m(x))^T$ — вектор базисных функций.

В случае, когда измерения по схеме (1.1) проводятся в точках $x_1, \dots, x_N \in X$ (некоторые из этих точек могут совпадать), получаем классическую линейную регрессионную модель $(Y, F\theta, \sigma^2 I_N)$, где

$$Y = \begin{bmatrix} y_1 \\ \vdots \\ y_N \end{bmatrix}, \quad F = \begin{bmatrix} f_1(x_1) & \dots & f_m(x_1) \\ \vdots & \dots & \vdots \\ f_1(x_N) & \dots & f_m(x_N) \end{bmatrix} = \begin{bmatrix} f^T(x_1) \\ \vdots \\ f^T(x_N) \end{bmatrix}. \quad (1.3)$$

Обычная запись из § 1 гл. 1 получается, если положить $x_{ji} = f_i(x_j)$ ($j = 1, \dots, N$; $i = 1, \dots, m$). Как известно из п. 1.2 гл. 1, другой формой записи модели $(Y, F\theta, \sigma^2 I_N)$ является

$$Y = F\theta + e, \quad e \sim (0, \sigma^2 I_N). \quad (1.4)$$

Как было показано в гл. 1, точность оценивания параметров θ зависит от матрицы F или (в нашем случае) от выбора точек x_j . Это дает возможность ставить задачу об оптимальном выборе точек — планировании эксперимента. Набор этих точек (возможно, повторяющихся) иногда называют *планом эксперимента* (дискретным, ненормированным).

Для нас удобнее называть планом эксперимента

$$\xi = \xi_N = \left\{ \begin{matrix} x_1 & \dots & x_N \\ 1/N & \dots & 1/N \end{matrix} \right\}, \quad (1.5)$$

т. е. вероятностную меру, сосредоточенную в этих точках с равными весами. Планы (1.5) назовем *дискретными (точными) планами*. Множество всевозможных дискретных планов вида (1.5) будем обозначать через Ξ_N .

В тех случаях, когда среди точек $x_1, \dots, x_N$ плана (1.5) имеются одинаковые, этот план удобно записывать в виде

$$\xi = \left\{ \begin{matrix} x_1 & \dots & x_n \\ p_1 & \dots & p_n \end{matrix} \right\}; \quad (1.6)$$

здесь $n \leq N$.

$$p_i \geq 0, \quad \sum_{i=1}^n p_i = 1, \quad p_i = r_i/N, \quad i = 1, \dots, n, \quad (1.7)$$

где r_i — число повторений i -й точки плана (1.6) в плане (1.5), соответствующее числу измерений, проводимых в точке x_i плана (1.6). План вида (1.6) с ограничениями (1.7) также является дискретным планом.

Одной из распространенных постановок задач планирования эксперимента по оцениванию параметров регрессии (1.2) является задача построения такого дискретного плана вида (1.5), что получаемые НЛН-оценки неизвестных параметров по результатам измерений (1.1) в точках этого плана были бы оптимальны в заданном смысле.

Как мы увидим ниже, критерии оптимальности являются сложными функциями точек плана. В большинстве случаев получить точное решение указанной задачи оптимального планирования в множестве Ξ_N очень трудно. Тем не менее, если множество допустимых планов расширить, удастся получить содержательные математические результаты. Если считать число точек n ограниченным, а N сколь угодно большим, то приходим к так называемым *непрерывным планам* (в литературе можно также встретить термин «распределение эксперимента»). В английской научной литературе термину «непрерывный план» соответствует термин «приближенный план» (approximate design), более точно отражающий смысл плана.

Непрерывным планом называется план вида (1.6), для которого ограничения (1.7) ослаблены: $p_i \geq 0$, $\sum_{i=1}^n p_i = 1$ ($i = 1, \dots, n$). Таким образом, непрерывным планом является вероятностная мера

$$\xi = \left\{ \begin{matrix} x_1, & \dots, & x_n \\ p_1, & \dots, & p_n \end{matrix} \right\}, \quad p_i \geq 0, \quad \sum_{i=1}^n p_i = 1, \quad (1.8)$$

где x_i ($i = 1, \dots, n$) — произвольные точки множества X . Если все точки x_j ($j = 1, \dots, n$) в (1.8) различны, то говорят, что план ξ сосредоточен в n точках.

В отличие от множества дискретных планов Ξ_N множество всевозможных непрерывных планов вида (1.8) является выпуклым, т. е. для любого $\alpha \in [0, 1]$ и любых планов ξ_1, ξ_2 вида (1.8) выпуклая комбинация этих планов $(1 - \alpha)\xi_1 + \alpha\xi_2$ может быть представлена в виде (1.8) и поэтому также является непрерывным планом. Действительно, поясним, что собой представляет выпуклая комбинация $\xi = (1 - \alpha)\xi_1 + \alpha\xi_2$ непрерывных планов ξ_1 и ξ_2 . Пусть точки x_i , на которых сосредоточены ξ_1 и ξ_2 , совпадают (если это не так, то несовпадающие точки

добавим в планы и припишем им нулевые веса), а сами эти планы имеют вид

$$\xi_1 = \left\{ \begin{matrix} x_1, & \dots, & x_n \\ p_1^{(1)}, & \dots, & p_n^{(1)} \end{matrix} \right\}, \quad \xi_2 = \left\{ \begin{matrix} x_1, & \dots, & x_n \\ p_1^{(2)}, & \dots, & p_n^{(2)} \end{matrix} \right\}.$$

Тогда

$$\xi = (1 - \alpha)\xi_1 + \alpha\xi_2 = \left\{ \begin{matrix} x_1, & \dots, & x_n \\ (1 - \alpha)p_1^{(1)} + \alpha p_1^{(2)}, & \dots, & (1 - \alpha)p_n^{(1)} + \alpha p_n^{(2)} \end{matrix} \right\},$$

т. е. ξ имеет вид (1.8) с $p_i = (1 - \alpha)p_i^{(1)} + \alpha p_i^{(2)}$ ($i = 1, \dots, n$).

С целью построения последовательной математической теории обобщим понятие непрерывного плана. *Непрерывным планом* (на множестве планирования X) будем называть произвольную вероятностную меру на измеримом пространстве $(X, \mathcal{B})$, где $\mathcal{B}$ есть σ -алгебра борелевских подмножеств множества X . Это обобщение позволяет привлечь для решения задач оптимального планирования удобный математический аппарат решения экстремальных задач в замкнутых множествах.

Множество всевозможных непрерывных планов на X будем обозначать через Ξ .

1.2. Информационная и дисперсионная матрицы. Информационной матрицей (Фишера) для модели (1.4) называется матрица $M = F^T F$. Поскольку для рассматриваемой схемы измерений и плана эксперимента (1.5) матрица F имеет вид (1.3), то

$$M = F^T F = (f(x_1), \dots, f(x_N)) \begin{bmatrix} f^T(x_1) \\ \vdots \\ f^T(x_N) \end{bmatrix} = \sum_{i=1}^N f(x_i) f^T(x_i). \quad (1.9)$$

Если план эксперимента имеет вид (1.6) и удовлетворяет ограничениям (1.7), то информационная матрица равна

$$M = \sum_{i=1}^n r_i f(x_i) f^T(x_i) = N \sum_{i=1}^n p_i f(x_i) f^T(x_i), \quad (1.10)$$

где (согласно (1.7)) $p_i = r_i/N$ ($i = 1, \dots, N$).

Информационная матрица (1.10) играет, как следует из результатов гл. 1, определяющую роль при построении наилучших линейных оценок неизвестных параметров θ модели (1.4). В частности, если матрица $F^T F$ невырожденная, то НЛН-оценка $\hat{\theta}$ параметров θ имеет вид

$$\hat{\theta} = (F^T F)^{-1} F^T Y, \quad (1.11)$$

а дисперсионная матрица этой оценки равна

$$D\hat{\theta} = \sigma^2 (F^T F)^{-1}. \quad (1.12)$$

При математических построениях, связанных с оптимальным планированием эксперимента, более удобно работать не с информационными матрицами вида (1.9) и (1.10), а с нормированными информационными матрицами.

Нормированной информационной матрицей дискретного плана ξ_N вида (1.5) называется

$$M(\xi_N) = \frac{1}{N} \sum_{i=1}^N f(x_i) f^T(x_i). \quad (1.13)$$

Если план ξ_N имеет вид (1.6) и удовлетворяет ограничениям (1.7), то

$$M(\xi_N) = \sum_{i=1}^n \frac{r_i}{N} f(x_i) f^T(x_i) = \sum_{i=1}^n p_i f(x_i) f^T(x_i).$$

По аналогии с этим *нормированной информационной матрицей непрерывного плана* (1.8) называется матрица

$$M(\xi) = \sum_{i=1}^n p_i f(x_i) f^T(x_i), \quad (1.14)$$

а произвольного плана $\xi = \xi(dx)$ — матрица

$$M(\xi) = \int_X f(x) f^T(x) \xi(dx), \quad (1.15)$$

где интегрирование ведется по вероятностной мере, являющейся планом эксперимента, а интеграл от матрицы — это по определению матрица из интегралов ее элементов. Другая форма записи матрицы (1.15):

$$M(\xi) = \left\| \int_X f_i(x) f_j(x) \xi(dx) \right\|_{i,j=1}^m.$$

Далее слово «нормированная» будет опускаться и, следовательно, под информационной матрицей плана будет пониматься одна из матриц (1.13)–(1.15).

Пусть задан дискретный план ξ_N вида (1.5) с невырожденной информационной матрицей (1.13). Выразим определяемую по (1.12) дисперсионную матрицу НЛН-оценок $\hat{\theta}$ параметров θ через (1.13):

$$D\hat{\theta} = \sigma^2 (F^T F)^{-1} = \sigma^2 \frac{1}{N} \left(\frac{1}{N} F^T F \right)^{-1} = \frac{\sigma^2}{N} (M(\xi_N))^{-1}.$$

Матрицу $D(\xi_N) = (M(\xi_N))^{-1}$ будем называть *дисперсионной матрицей дискретного плана* ξ_N . Аналогично, если ξ — произвольный непрерывный план с невырожденной информационной матрицей $M(\xi)$, то *дисперсионной матрицей плана* ξ будем называть матрицу

$$D(\xi) = (M(\xi))^{-1}.$$

В дальнейшем планы, информационные матрицы которых невырождены, будем называть *невырожденными планами*.

Отметим, что как информационная, так и дисперсионная матрицы плана не зависят от неизвестных параметров (включая

и σ^2) и результатов измерений. Кроме того, если имеется дискретный план ξ вида (1.6) с ограничениями (1.7), то обе матрицы $M(\xi)$ и $D(\xi)$ не зависят от того, проводить по этому плану N или kN измерений (при любом натуральном k), выбирая точки x_i по kr_i раз ($i = 1, \dots, n$).

Используя указанные свойства информационных и дисперсионных матриц, качество непрерывных планов будем выражать через эти матрицы. Если имеются такие два невырожденных плана ξ_1 и ξ_2 , что

$$D(\xi_2) < D(\xi_1), \quad (1.16)$$

будем считать, что план ξ_2 лучше плана ξ_1 . Как правило, будем считать также, что план ξ_2 лучше плана ξ_1 , если вместо строгого неравенства (1.16) выполнено нестрогое $D(\xi_2) \leq D(\xi_1)$ (т. е. матрица $D(\xi_1) - D(\xi_2)$ неотрицательно определена) и $D(\xi_1) \neq D(\xi_2)$.

1.3. Примеры сравнения планов. Приведем два простых примера.

Пример 1.1. Дана линейная регрессия на отрезке $X = [-1, 1]$:

$$y_i = \theta_1 + \theta_2 x_i + \varepsilon_i, \quad E\varepsilon_i = 0, \quad E\varepsilon_i \varepsilon_j = 0, \quad i \neq j, \quad E\varepsilon_i^2 = \sigma^2. \quad (1.17)$$

Здесь $\theta = (\theta_1, \theta_2)^T$, $f(x) = (1, x)^T$.

В качестве плана ξ_1 выберем план

$$\xi_1 = \left\{ \begin{array}{ccc} -1 & 0 & 1 \\ 1/3 & 1/3 & 1/3 \end{array} \right\}.$$

сосредоточенный в точках $x_1 = -1$, $x_2 = 0$, $x_3 = 1$. Для этого плана

$$M(\xi_1) = \frac{1}{3} \sum_{i=1}^3 \begin{bmatrix} 1 \\ x_i \end{bmatrix} (1, x_i) = \frac{1}{3} \sum_{i=1}^3 \begin{bmatrix} 1 & x_i \\ x_i & x_i^2 \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 2/3 \end{bmatrix},$$

$$D(\xi_1) = \begin{bmatrix} 1 & 0 \\ 0 & 3/2 \end{bmatrix}.$$

План ξ_2 выберем в виде

$$\xi_2 = \left\{ \begin{array}{cc} -1 & 1 \\ 1/2 & 1/2 \end{array} \right\}$$

(т. е. план ξ_2 сосредоточен на концах интервала $[-1, 1]$ с равными весами). Имеем

$$M(\xi_2) = \frac{1}{2} \begin{bmatrix} 1 & 1 \\ 1 & 1 \end{bmatrix} + \frac{1}{2} \begin{bmatrix} 1 & -1 \\ -1 & 1 \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix},$$

$$D(\xi_2) = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \leq D(\xi_1).$$

Пусть число измерений, проводимых по планам ξ_1 и ξ_2 , кратно 6, т. е. $N = 6k$. Тогда при использовании плана ξ_1 НЛН-оценок $\hat{\theta}_1$, $\hat{\theta}_2$ параметров θ_1 , θ_2 некоррелированы, а их дисперсии равны $D\hat{\theta}_1 = \sigma^2/(6k)$, $D\hat{\theta}_2 = \sigma^2/(4k)$. При использовании плана ξ_2

эти оценки также некоррелированы, а их дисперсии равны $D\hat{\theta}_1 = \sigma^2/(6k)$ и $D\hat{\theta}_2 = \sigma^2/(6k)$. Следовательно, план ξ_2 лучше плана ξ_1 (даже несмотря на то, что строгое неравенство (1.16) не выполнено).

Пример 1.2. Рассмотрим введенную в примере 1.4 гл. 1 схему взвешивания трех предметов на двухчашечных весах:

$$y_i = \theta_1 x_{i1} + \theta_2 x_{i2} + \theta_3 x_{i3} + \varepsilon_i, \\ E\varepsilon_i = 0, \quad D\varepsilon_i = \sigma^2, \quad E\varepsilon_i \varepsilon_j = 0, \quad i \neq j,$$

где x_{ij} могут принимать лишь три значения, а именно $-1, 1, 0$, в зависимости соответственно от того, взвешивался i -й предмет в j -м измерении на правой или левой чашке или вообще не взвешивался.

Пусть сначала число измерений равно 24, а план ξ_1 сосредоточен в точках

$$(1, 0, 0)^T, \quad (0, 1, 0)^T, \quad (0, 0, 1)^T$$

с равными весами (это означает, что по плану ξ_1 взвешивается каждый предмет отдельно на левой чашке по восемь раз). Вес каждого предмета оценивается по МНК:

$$\hat{\theta}_l = \frac{1}{8} \sum_{i=1}^N y_i x_{il}, \quad l = 1, 2, 3.$$

При этом оценки некоррелированы и $D\hat{\theta}_l = \sigma^2/8$, поскольку

$$M(\xi_1) = \begin{bmatrix} 1/3 & 0 & 0 \\ 0 & 1/3 & 0 \\ 0 & 0 & 1/3 \end{bmatrix}, \quad D(\xi_1) = \begin{bmatrix} 3 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & 3 \end{bmatrix}.$$

Положим теперь $N=8$ и план ξ_2 запишем в виде (1.5), где точки плана — различные комбинации $(\pm 1, \pm 1, \pm 1)$. Для такого плана НЛН-оценки параметров θ_l ($l=1, 2, 3$) некоррелированы и равны

$$\hat{\theta}_l = \frac{1}{8} \sum_{i=1}^8 y_i x_{il}, \quad D\hat{\theta}_l = \frac{\sigma^2}{8}$$

(в данном случае матрицы $M(\xi_2)$ и $D(\xi_2)$ единичные). Таким образом, проводя восемь измерений по плану ξ_2 , мы достигаем той же точности МНК-оценок неизвестных параметров, что и после 24 измерений по плану ξ_1 (который на первый взгляд кажется подходящим). Ясно, что в данном случае выполнено строгое неравенство (1.16).

1.4. Свойства информационных матриц. Согласно (1.15), информационная матрица плана ξ записывается в виде

$$M(\xi) = \int_{\mathbf{x}} f(\mathbf{x}) f^T(\mathbf{x}) \xi(d\mathbf{x}).$$

Этот вид позволяет вывести полезные свойства информационных матриц, которые сформулированы ниже в виде отдельных утверждений.

Лемма 1.1 (свойство 1). *Для любого плана ξ информационная матрица $M(\xi)$ неотрицательно определена (в частности, является симметричной).*

Доказательство легко следует из определения неотрицательной определенности матриц и того, что для всех $z \in \mathbb{R}^m$ выполнено

$$z^T M(\xi) z = \int_X [z^T f(x)]^2 \xi(dx) \geq 0.$$

Лемма 1.2 (свойство 2). *Если план ξ имеет вид (1.8) и $n < m$, то этот план вырожденный, т. е. $\det M(\xi) = 0$.*

Доказательство. В силу (1.14)

$$M(\xi) = \sum_{i=1}^n \rho_i f(x_i) f^T(x_i).$$

Используя теорему 1.13 из приложения 1, получаем

$$\text{rang } M(\xi) \leq \sum_{i=1}^n \text{rang } f(x_i) f^T(x_i).$$

Утверждение леммы будет доказано, если показать, что $\text{rang } f(x) f^T(x) \leq 1$ для любого фиксированного $x \in X$.

Выберем такое $j \in \{1, \dots, m\}$, что $f_j(x) \neq 0$ (если такого j не существует, то $\text{rang } f(x) f^T(x) = 0$). Каждый l -й столбец матрицы

$$f(x) f^T(x) = \|f_l(x) f_l(x)\|_{l, l-1}^m$$

можно представить в виде произведения j -го столбца на $f_l(x)/f_j(x)$.

Следовательно, все столбцы матрицы $f(x) f^T(x)$ линейно зависят от одного из них; поэтому $\text{rang } f(x) f^T(x) = 1$. Лемма доказана.

Обозначим через $\mathfrak{M}$ множество информационных матриц $M(\xi)$, соответствующее всевозможным непрерывным планам $\xi \in \Xi$.

Лемма 1.3 (свойство 3). *Множество $\mathfrak{M}$ выпукло*).*

Доказательство. Необходимо проверить, что для любого $\alpha \in [0, 1]$ и для любых непрерывных планов ξ_1, ξ_2 матрица

$$M = (1 - \alpha) M(\xi_1) + \alpha M(\xi_2) \quad (1.18)$$

принадлежит множеству $\mathfrak{M}$. Определим план ξ по формуле

$$\xi = (1 - \alpha) \xi_1 + \alpha \xi_2.$$

*) Множество V называется выпуклым, если вместе с любыми двумя точками v_1 и v_2 оно содержит и весь отрезок $\{v | v = \alpha v_1 + (1 - \alpha) v_2, \alpha \in [0, 1]\}$.

Покажем, что $M = M(\xi)$; действительно,

$$M = (1 - \alpha) M(\xi_1) + \alpha M(\xi_2) = (1 - \alpha) \int f(x) f^T(x) \xi_1(dx) + \\ + \alpha \int f(x) f^T(x) \xi_2(dx) = \int f(x) f^T(x) [(1 - \alpha) \xi_1(dx) + \alpha \xi_2(dx)] = M(\xi).$$

Лемма доказана.

Кроме выпуклости важным с математической точки зрения является свойство компактности множества информационных матриц $\mathfrak{M}$. Для полного понимания доказательства этого свойства следует ознакомиться с понятием слабой сходимости вероятностных мер и ее основными свойствами [45]. Нам потребуются определение и одно фундаментальное свойство. Напомним их.

Пусть вероятностные меры P_i ($i = 1, 2, \dots$) и P заданы на борелевских подмножествах множества X . Говорят, что последовательность $\{P_i\}_{i=1}^{\infty}$ слабо сходится к P , если для любой непрерывной ограниченной функции g на X выполнено условие

$$\int_X g(x) P_i(dx) \rightarrow \int_X g(x) P(dx) \quad (i \rightarrow \infty).$$

Лемма 1.4. *Если множество X компактно, то семейство вероятностных мер, заданных на борелевских подмножествах множества X , слабо компактно (т. е. из любой последовательности вероятностных мер можно выделить подпоследовательность, слабо сходящуюся к некоторой вероятностной мере, определенной там же).*

Доказательство имеется в [45].

В случае $X \subseteq \mathbb{R}$ утверждение леммы 1.4 называется *теоремой Хелли*.

Лемма 1.5 (свойство 4). *Если множество планирования X — компакт, а функции $f_i(x)$ ($i = 1, \dots, m$) непрерывны на X , то множество информационных матриц $\mathfrak{M}$ компактно.*

Доказательство. В силу леммы 1.4 множество непрерывных планов Ξ слабо компактно. Следовательно, из каждой последовательности планов $\{\xi_i\}_{i=1}^{\infty}$ можно выделить слабо сходящуюся подпоследовательность $\{\xi_{i_j}\}_{j=1}^{\infty}$, причем предел ξ_* указанной подпоследовательности в силу компактности X является вероятностной мерой, т. е. непрерывным планом. Слабая сходимость $\{\xi_{i_j}\}$ к ξ_* означает, что для любой непрерывной на X функции $g(x)$ имеет место равенство

$$\lim_{j \rightarrow \infty} \int_X g(x) \xi_{i_j}(dx) = \int_X g(x) \xi_*(dx).$$

Выбирая в качестве $g(x)$ компоненты матрицы $f(x)f^T(x)$, получаем

$$\lim_{j \rightarrow \infty} M(\xi_{i_j}) = M(\xi_*).$$

Отсюда следует, что из любой последовательности информационных матриц $\{M(\xi_i)\}$ можно выделить подпоследовательность, сходящуюся к элементу множества $\mathcal{M}$. Это и означает, что множество $\mathcal{M}$ компактно. Лемма доказана.

Для вывода еще одного свойства информационных матриц потребуются приведенные ниже сведения из выпуклого анализа.

Выпуклой оболочкой множества V называется пересечение всех выпуклых множеств, содержащих множество V . Выпуклая оболочка множества V обозначается символом $\text{conv } V$.

Пусть $v_1, \dots, v_k$ — точки из некоторого множества V . Выпуклой комбинацией точек $v_1, \dots, v_k$ называется точка

$$v = \sum_{i=1}^k \lambda_i v_i$$

где

$$\lambda_i \geq 0, \quad i = 1, \dots, k, \quad \sum_{i=1}^k \lambda_i = 1.$$

Точка v множества V называется крайней точкой этого множества, если она не представима в виде выпуклой комбинации любых других точек множества V , т.е. если из того, что

$$v = (1 - \alpha) v_1 + \alpha v_2, \quad 0 < \alpha < 1, \quad v_1, v_2 \in V,$$

следует $v_1 = v_2$.

Докажем простое вспомогательное утверждение.

Лемма 1.6. *Выпуклая оболочка множества V состоит из всевозможных выпуклых комбинаций точек множества V .*

Доказательство. Обозначим через U множество всех выпуклых комбинаций точек из V . Имеем $V \subset U$. Множество U выпукло, поскольку если

$$v_1 = \sum_{i=1}^{m_1} \alpha_{i1} v_{i1}, \quad v_2 = \sum_{i=1}^{m_2} \alpha_{i2} v_{i2},$$

$$\sum_{i=1}^{m_l} \alpha_{il} = 1, \quad \alpha_{il} \geq 0, \quad i = 1, \dots, m_l, \quad l = 1, 2,$$

то для любого $\alpha \in [0, 1]$

$$\alpha v_1 + (1 - \alpha) v_2 = \sum_{i=1}^{m_1} \alpha \alpha_{i1} v_{i1} + \sum_{i=1}^{m_2} (1 - \alpha) \alpha_{i2} v_{i2}$$

и, следовательно, $\alpha v_1 + (1 - \alpha) v_2$ является выпуклой комбинацией точек $v_{11}, v_{21}, \dots, v_{m_1 1}, v_{12}, \dots, v_{m_2 2}$. Отсюда получаем, что $\text{conv } V \subset U$. С другой стороны, по индукции доказывается, что каждая точка из U содержится в любом выпуклом множестве, содержащем V , и поэтому $U \subset \text{conv } V$. Лемма доказана.

Лемма 1.7. *Пусть $V \subset R^k$, $v_i \in V$ ($i = 1, \dots, r$). Если $r \geq k + 2$, то существуют такие не равные одновременно нулю числа $\beta_1, \dots, \beta_r$, что*

$$\sum_{i=1}^r \beta_i = 0, \quad \sum_{i=1}^r \beta_i v_i = 0. \quad (1.19)$$

Доказательство. Точке

$$v_i = (v_i(1), \dots, v_i(k))^T \in R^k$$

поставим в соответствие точку

$$\bar{v}_i = (1, v_i(1), \dots, v_i(k))^T \in \mathbb{R}^{k+1}.$$

В $\mathbb{R}^{k+1}$ любые $k+2$ векторов линейно зависимы, и поэтому существуют такие не равные одновременно нулю числа $\beta_1, \dots, \beta_r$, что

$$\sum_{i=1}^r \beta_i \bar{v}_i = 0.$$

Приравнивая нулю отдельно первую координату и n оставшихся, получаем требуемое. Лемма доказана.

Отметим, что если система уравнений (1.19) имеет только нулевое решение $\beta_1 = \dots = \beta_r = 0$, то множество векторов $\{v_1, \dots, v_r\}$ называется аффинно зависимым.

В лемме 1.7 утверждается таким образом, что в $\mathbb{R}^k$ существует не более чем $k+1$ аффинно независимых векторов.

Теорема 1.1 (теорема Каратеодори). Пусть V — компактное подмножество $\mathbb{R}^n$. Тогда каждая точка множества $\text{conv } V$ может быть представлена в виде выпуклой комбинации с положительными коэффициентами не более чем $k+1$ крайних точек множества V .

Доказательство. В силу леммы 1.6 любая точка v из $\text{conv } V$ может быть представлена в виде

$$v = \sum_{i=1}^r \alpha_i v_i, \quad \alpha_i > 0, \quad v_i \in V, \quad i = 1, \dots, r, \quad \sum_{i=1}^r \alpha_i = 1$$

(в определении выпуклой комбинации $\alpha_i \geq 0$, но нулевые α_i можно опустить). Допустим, что $r \geq n+2$. В силу леммы 1.7 существуют такие не равные одновременно нулю числа $\beta_1, \dots, \beta_r$, что выполнено (1.19). Положим

$$\varepsilon = \min_{\{i | \beta_i > 0\}} (\alpha_i / \beta_i), \quad \bar{\alpha}_i = \alpha_i - \varepsilon \beta_i, \quad i = 1, \dots, r.$$

Поскольку среди коэффициентов $\beta_1, \dots, \beta_r$ существуют положительные, то $\varepsilon > 0$. Величины $\bar{\alpha}_i$ удовлетворяют соотношениям

$$\sum_{i=1}^r \bar{\alpha}_i = 1, \quad \bar{\alpha}_i \geq 0, \quad i = 1, \dots, r.$$

причем по крайней мере одно из чисел $\bar{\alpha}_i$ обращается в нуль. Замечая теперь, что

$$v = \sum_{i=1}^r \bar{\alpha}_i v_i,$$

получаем представление точки v как выпуклой комбинации $r-1$ точек из V . Продолжая далее аналогично, получаем представление, в котором $r \leq k+1$. Теорема доказана.

Следствие 1.1 (свойство 5). Если множество $\mathfrak{M}$ компактно, то для любого плана ξ информационная матрица $M(\xi)$ представима в виде

$$M(\xi) = \sum_{i=1}^{n_0} p_i M(x_i), \quad 0 \leq p_i \leq 1, \quad \sum_{i=1}^{n_0} p_i = 1,$$

где $M(x) = f(x)f^T(x)$ — информационная матрица плана ξ , сосредоточенного в точке x , и

$$n_0 \leq \frac{1}{2} m(m+1) + 1. \quad (1.20)$$

Доказательство. Так как любая информационная матрица симметрична, то она полностью определяется своими $l = m(m+1)/2$ элементами, т. е. ей можно сопоставить вектор размерности l . Утверждение следствия вытекает из теоремы Каратеодори и того очевидного факта, что крайними в множестве $\mathcal{M}$ являются только точки вида $M(\xi_x) = f(x)f^T(x)$, соответствующие планам $\xi_x = \begin{Bmatrix} x \\ 1 \end{Bmatrix}$. Следствие доказано.

Из следствия 1.1 вытекает, что непрерывные планы с требуемыми свойствами информационной матрицы (в частности, оптимальные планы) можно строить на множестве планов вида (1.8) с $n \leq m(m+1)/2 + 1$.

1.5. Критерии оптимальности. Как указано в п. 1.2, матрица $D(\xi) = M^{-1}(\xi)$ является матричной характеристикой невырожденного плана ξ .

Приведем пример, показывающий, что не из любой пары планов ξ_1, ξ_2 можно выбрать тот, который лучше.

Пример 1.3. Пусть функция регрессии имеет вид $\eta(x) = \theta_1 + \theta_2 x + \theta_3 x^2$, $x \in [-1, 1]$. Здесь $f(x) = (1, x, x^2)$. Положим

$$\xi_1 = \begin{Bmatrix} -1 & 0 & 1 \\ 1/3 & 1/3 & 1/3 \end{Bmatrix}, \quad \xi_2 = \begin{Bmatrix} -1 & 0 & 1 \\ 1/4 & 1/2 & 1/4 \end{Bmatrix}.$$

Имеем

$$M(\xi_1) = \sum_{i=1}^3 \frac{1}{3} \begin{bmatrix} 1 & x_i & x_i^2 \\ x_i & x_i^2 & x_i^3 \\ x_i^2 & x_i^3 & x_i^4 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 2/3 \\ 0 & 2/3 & 0 \\ 2/3 & 0 & 2/3 \end{bmatrix},$$

$$\det M(\xi_1) = \frac{4}{27},$$

$$D(\xi_1) = \begin{bmatrix} 3 & 0 & -3 \\ 0 & 3/2 & 0 \\ -3 & 0 & 9/2 \end{bmatrix},$$

$$M(\xi_2) = \begin{bmatrix} 1 & 0 & 1/2 \\ 0 & 1/2 & 0 \\ 1/2 & 0 & 1/2 \end{bmatrix}, \quad \det M(\xi_2) = \frac{1}{8}, \quad D(\xi_2) = \begin{bmatrix} 2 & 0 & -2 \\ 0 & 2 & 0 \\ -2 & 0 & 4 \end{bmatrix}.$$

Матрица

$$D(\xi_1) - D(\xi_2) = \begin{bmatrix} 1 & 0 & -1 \\ 0 & -1/2 & 0 \\ -1 & 0 & 1/2 \end{bmatrix}$$

имеет на главной диагонали как положительные, так и отрицательные элементы и поэтому не свидетельствует о том, что какой-либо из планов безусловно лучше другого. Отметим, что оба плана (ξ_1 и ξ_2) оптимальны, но по разным критериям.

Поскольку в подавляющем большинстве случаев задача минимизации матрицы $D(\xi)$ однозначного решения не имеет, то рассматривают задачу минимизации некоторого функционала Φ ,

заданного на множестве дисперсионных матриц, т. е. задачу поиска плана

$$\xi^* = \arg \min \Phi [D(\xi)] \quad (1.21)$$

на фиксированном множестве планов (Ξ или Ξ_N). При этом, не оговаривая особо, всегда предполагается, что

$$\min_{\xi \in \Xi} \Phi [D(\xi)] > -\infty.$$

План (1.21) называется *Φ-оптимальным* (или просто *оптимальным*). Функционал Φ называется *критерием оптимальности*. Если допустимое множество планов — Ξ (что далее почти всегда и предполагается), то задача построения плана (1.21) называется задачей *непрерывного оптимального планирования* (или просто *оптимального планирования*).

Основными свойствами, которыми обладают практически все функционалы Φ , имеющие статистический смысл, являются следующие:

- а) монотонность $\Phi(D_1) \leq \Phi(D_2)$, если $D_1 \leq D_2$;
- б) однородность $\Phi(\beta D) = \gamma(\beta)\Phi(D)$,

где $\gamma(\beta)$ — возрастающая функция;

- в) выпуклость

$$\Phi[D((1-\alpha)\xi_1 + \alpha\xi_2)] \leq (1-\alpha)\Phi[D(\xi_1)] + \alpha\Phi[D(\xi_2)] \quad (1.22)$$

для любых планов $\xi_1, \xi_2 \in \Xi$ и $0 < \alpha < 1$.

Большинство приведенных ниже функционалов строго выпуклые, т. е. для них в (1.22) имеет место строгое неравенство для всех $0 < \alpha < 1$ при $D(\xi_1) \neq D(\xi_2)$. Отметим, что условие однородности обеспечивает независимость оптимального плана от N .

Наиболее нетривиальным из выписанных условий является условие выпуклости (и строгой выпуклости). Доказательства выпуклости приводимых ниже функционалов содержатся в § 2 приложения 1; результаты о выпуклости очень важны для понимания материала данного пункта.

Приведем два почти очевидных свойства непрерывных оптимальных планов.

Теорема 1.2. 1) Любая комбинация непрерывных оптимальных планов является непрерывным оптимальным планом, т. е. множество непрерывных оптимальных планов выпукло.

2) Если функционал Φ строго выпуклый, то все непрерывные оптимальные планы имеют одну и ту же дисперсионную матрицу.

Доказательство. 1) Пусть планы ξ_1 и ξ_2 оптимальны. Тогда из (1.22) следует

$$\Phi[D(\xi^*)] \leq (1-\alpha)\Phi[D(\xi_1)] + \alpha\Phi[D(\xi_2)] = \min \Phi[D(\xi)],$$

где $\xi^* = (1 - \alpha)\xi_1^* + \alpha\xi_2^*$. Следовательно,

$$\Phi[D(\xi^*)] = \min_{\xi} \Phi[D(\xi)],$$

и поэтому план ξ^* оптимален.

2) Предположим, что имеется два оптимальных плана ξ_1^* , ξ_2^* , причем $D(\xi_1^*) \neq D(\xi_2^*)$. В силу строгой выпуклости Φ для плана

$$\xi^* = (1 - \alpha)\xi_1^* + \alpha\xi_2^*, \quad 0 < \alpha < 1,$$

имеем

$$\Phi[D(\xi^*)] < (1 - \alpha)\Phi[D(\xi_1^*)] + \alpha\Phi[D(\xi_2^*)] = \min_{\xi} \Phi[D(\xi)],$$

но в силу 1) план ξ^* должен быть оптимальным. Теорема доказана.

Для того чтобы вывести еще одно, более нетривиальное свойство оптимальных планов, докажем следующую теорему, являющуюся дополнением к теореме Каратеодори (теореме 1.1).

Теорема 1.3. Пусть V — компактное подмножество $\mathbb{R}^k$, v — граничная точка множества $\text{conv } V$. Тогда v может быть представлена в виде выпуклой комбинации с положительными коэффициентами не более чем k точек из V .

Доказательство.* В силу теоремы 1.1 существуют такие векторы $v_1, \dots, v_{k+1}$ из V , что

$$v = \sum_{i=1}^{k+1} \alpha_i v_i, \quad \sum_{i=1}^{k+1} \alpha_i = 1. \quad (1.23)$$

причем будем считать, что все α_i ($i = 1, \dots, k+1$) положительны, а векторы $v_1, \dots, v_{k+1}$ аффинно независимы (иначе доказывать нечего). Поскольку v — граничная точка выпуклого компакта $\text{conv } V$, то по теореме об отделимости (см. [2, с. 53]) найдется такой ненулевой вектор $a = (a_1, \dots, a_k)^T$, что

$$(a, u) \leq (a, v)$$

для всех $u \in \text{conv } V$ (здесь $(\cdot, \cdot)$ — скалярное произведение в $\mathbb{R}^k$). В частности,

$$(a, v_i) \leq (a, v) \quad (1.24)$$

при всех $i = 1, \dots, k+1$.

Из (1.23) имеем

$$(a, v) = \sum_{i=1}^{k+1} \alpha_i (a, v_i), \quad \sum_{i=1}^{k+1} \alpha_i = 1.$$

Отсюда с учетом положительности всех α_i и неравенств (1.24) следует, что последние являются равенствами. Положим

$$a_{k+1} = -(a, v) = -(a, v_i), \quad i = 1, \dots, k+1.$$

Относительно неизвестных коэффициентов $a_1, \dots, a_{k+1}$ имеем, следовательно, систему уравнений

$$(a, v_i) + a_{k+1} = 0, \quad i = 1, \dots, k+1 \quad (1.25)$$

* Приводимое доказательство принадлежит В. Н. Малоземову.

Аффинная независимость векторов $v_1, \dots, v_{n+1}$ эквивалентна, как нетрудно понять, тому, что система уравнений (1.25) не имеет нетривиальных решений. Следовательно, $a = 0$, и поэтому получили противоречие. Теорема доказана.

Следующая теорема дополняет следствие 1.1.

Теорема 1.4. Пусть множество $\mathfrak{M}$ компактно, выполнено условие однородности и оптимальный план ξ^* невырожден. Тогда этот план может быть записан в виде (1.8) с $n \leq m(m+1)/2$.

Доказательство. В силу теоремы 1.3 достаточно показать, что матрица $M(\xi^*)$ принадлежит границе множества $\mathfrak{M}$. Действительно, если $M(\xi^*)$ — внутренняя точка $\mathfrak{M}$, то существует такое $\alpha > 1$, что $\alpha M(\xi^*) \in \mathfrak{M}$, причем $\alpha M(\xi^*)$ — информационная матрица некоторого плана ξ . Поэтому

$$\Phi[M^{-1}(\xi^*)] > \gamma\left(\frac{1}{\alpha}\right) \Phi[M^{-1}(\xi^*)] = \Phi[\alpha^{-1}M^{-1}(\xi^*)] = \Phi[M^{-1}(\xi)],$$

что противоречит определению оптимального плана. Теорема доказана.

С информационными матрицами работать проще, чем с дисперсионными, поэтому критерий оптимальности часто записывают как функционал Ψ , заданный на множестве информационных матриц $\mathfrak{M}$.

План

$$\xi^* = \arg \min \Psi[M(\xi)] \quad (1.26)$$

называют Ψ -оптимальным.

Задачу Ψ -оптимального планирования (1.26) рассматривают для случая, когда функционал Ψ выпуклый (т. е. удовлетворяющий неравенству (1.22) с заменой Φ, D на Ψ, M), монотонный и однородный.

Известно довольно много критериев оптимальности планов. Ниже приведены наиболее часто используемые из них.

План

$$\xi^* = \arg \max_{\xi} \det M(\xi) = \arg \min_{\xi} \det D(\xi) \quad (1.27)$$

называется D -оптимальным. Функционалы

$$\Phi[D] = \ln \det D, \quad \Psi[M] = -\ln \det M \quad (1.28)$$

называются *критериями D -оптимальности* (от английского термина «determinant», т. е. «определитель»).

Строгая выпуклость функционалов Φ и Ψ , определяемых по формуле (1.28), следует из формулы (2.6) приложения 1. Для функционала $\Phi[D] = \ln \det D$ не выполнено условие однородности $\Phi[\alpha D] = \gamma(\alpha)\Phi[D]$, но выполнено условие

$$\Phi[\alpha D] = m \ln \alpha + \ln \det D, \quad (1.29)$$

которое его заменяет. Отметим, что утверждение теоремы 1.4 остается справедливым при замене условия однородности на (1.29).

При D оптимальном планировании минимизируется объем области, ограниченной эллипсоидом рассеяния НЛН-оценок $\hat{\theta}$ неизвестных параметров.

Напомним (более подробно это понятие изучается в курсе математической статистики), что если имеется m -мерный случайный вектор η с вектором средних $E\eta$ и дисперсионной матрицей $D\eta$, то эллипсоид рассеяния этого случайного вектора определяется как эллипсоид, равномерное распределение внутри которого имеет вектор средних $E\eta$ и дисперсионную матрицу $D\eta$. Этот эллипсоид имеет вид

$$\{x \in \mathbb{R}^m \mid (x - E\eta)^T D\eta (x - E\eta) = m + 2\}.$$

Объем области, ограниченной этим эллипсоидом, как нетрудно убедиться прямым подсчетом m -мерного интеграла, равен

$$V = (m + 2)^{m+2} \pi^{m/2} (\det D\eta)^{1/2} / \Gamma\left(\frac{m}{2} + 1\right).$$

Для рассматриваемого случая, когда случайный вектор η представляет собой НЛН-оценку $\hat{\theta}$ параметров θ , эллипсоид рассеяния имеет вид

$$\{x \in \mathbb{R}^m \mid (x - \hat{\theta})^T D\hat{\theta} (x - \hat{\theta}) = m + 2\}, \quad (1.30)$$

а его объем пропорционален квадратному корню из определителя дисперсионной матрицы плана. Отметим также, что если распределение случайных ошибок e_i в (1.1) нормально, то распределение случайного вектора $\hat{\theta}$ также нормально, а эллипсоид рассеяния (1.30) является поверхностью равного уровня распределения случайного вектора $\hat{\theta}$.

Важным с практической точки зрения свойством непрерывных D -оптимальных планов является их инвариантность относительно любых невырожденных линейных преобразований базисных функций и неизвестных параметров.

Теорема 1.5. Пусть ξ^* есть D -оптимальный план для модели (1.4), L — произвольная невырожденная $m \times m$ -матрица. Тогда план ξ^* будет D -оптимальным, во-первых, для модели

$$(Y, \Phi\theta, \sigma^2 I_N), \quad \Phi = \begin{bmatrix} \varphi_1(x_1) & \dots & \varphi_m(x_1) \\ \vdots & \ddots & \vdots \\ \varphi_1(x_N) & \dots & \varphi_m(x_N) \end{bmatrix}, \quad \varphi(x) = Lf(x),$$

и, во-вторых, для модели

$$(Y, F\theta^{(1)}, \sigma^2 I_N), \quad \theta^{(1)} = L^T \theta.$$

Доказательство. Докажем первое из утверждений теоремы. Функцию регрессии запишем в виде

$$\eta(x, \theta) = \theta^T f(x) = \theta_{11}^T \varphi(x), \quad \theta_{11} = (L^{-1})^T \theta.$$

Пусть $\hat{\theta}$ — НЛН-оценка θ при любом плане ξ с дисперсионной матрицей $D(\hat{\theta}, \xi)$. Тогда, согласно теореме 1.4 из гл. 1,

$\hat{\theta}_{(1)} = (L^{-1})^T \hat{\theta}$ — НЛН-оценка параметров $\theta_{(1)}$, причем дисперсионная матрица этой оценки по лемме 13 из гл. 1 равна

$$D(\hat{\theta}_{(1)}, \xi) = (L^{-1})^T D(\hat{\theta}, \xi) L^{-1}.$$

Отсюда

$$\det D(\hat{\theta}_{(1)}, \xi) = (\det L)^{-2} \det D(\hat{\theta}, \xi).$$

Минимизируя обе части последнего неравенства по ξ , получаем

$$\xi^* = \arg \min_{\xi} \det D(\hat{\theta}, \xi) = \arg \min_{\xi} \det D(\hat{\theta}_{(1)}, \xi).$$

Второе утверждение теоремы доказывается аналогично. Теорема доказана.

Перейдем к следующему критерию оптимальности.

План

$$\xi^* = \arg \min_{\xi} \det [A^T D(\xi) A] \quad (1.31)$$

называется *обобщенно D-оптимальным*. Здесь A — матрица полного ранга размера $m \times s$, $s \leq m$. В наиболее важном частном случае $A = (I_s, 0)^T$ имеем

$$A^T D(\xi) A = D_s(\xi),$$

где $D_s(\xi)$ — главный минор порядка s матрицы $D(\xi)$, а план (1.31) называется *D_s -оптимальным*. Этот случай соответствует ситуации, когда экспериментатора интересуют только первые s из m неизвестных параметров. Отметим, что если $s = m$ (а поэтому и $\text{rang } A = m$), то обобщенно D -оптимальные планы совпадают с D -оптимальными.

Функционал

$$\Phi(D) = \ln \det (A^T D A) \quad (1.32)$$

называется *критерием обобщенной D-оптимальности*.

Выпуклость функционала (1.32), так же как и выпуклость D -критерия, следует из формулы (2.6) приложения 1.

Очевидна статистическая интерпретация D_s -оптимальности: для D_s -оптимального плана минимален объем эллипсоида рассеяния, соответствующего НЛН-оценкам интересующих нас s параметров.

Свойства критерия обобщенной D -оптимальности аналогичны свойствам D -критерия.

Критерий обобщенной D -оптимальности может быть использован в случае, когда необходимо оценить некоторую векторную параметрическую функцию $T\theta$ (T есть $s \times m$ -матрица). Для этого нужно сделать невырожденную замену параметров

$$\theta^{(1)} = A\theta = (T\theta, S\theta)^T$$

(где S есть $(m - s) \times m$ -матрица), а модель (1.4) заменить на

$$(Y, F^{(1)}(\theta^{(1)}), \sigma^2 I_N), \quad F^{(1)} = FA^{-1}.$$

Такая замена позволяет, в частности, повысить точность оценивания тех параметров (или их линейных комбинаций), которые особенно интересуют исследователя, за счет отказа от точности оценивания остальных.

Кроме D -критерия широко распространены также линейные критерии оптимальности.

Линейно оптимальным (или *L -оптимальным*) называется план

$$\xi^* = \arg \max_{\xi} \operatorname{tr} LD(\xi), \quad (1.33)$$

где L — фиксированная, неотрицательно определенная $m \times m$ -матрица. Функционал

$$\Phi(D) = \operatorname{tr} LD \quad (1.34)$$

называется *линейным критерием оптимальности* или *критерием L -оптимальности* (от английского слова «linear» — линейный).

Функционал (1.34) называется линейным потому, что для любых $m \times m$ -матриц A, B и положительного числа c имеют место равенства

$$\Phi(A + B) = \Phi(A) + \Phi(B), \quad \Phi(cA) = c\Phi(A).$$

Функционал Φ , не представимый в виде (1.34), указанными свойствами линейности обладать не может.

Выпуклость функционала (1.34) следует из теоремы 2.4 предложения 1, отсюда же следует строгая выпуклость (1.34) в случае $L > 0$.

Статистический смысл L -оптимального плана заключается в том, что для него обобщенные квадратичные потери

$$r(L, \bar{\theta}, \theta) = E(\bar{\theta} - \theta)^T L(\bar{\theta} - \theta)$$

минимальны (здесь $\bar{\theta} = \bar{\theta}(\xi)$ — НЛН-оценки для θ , соответствующие плану ξ). Это станет очевидным, если $r(L, \bar{\theta}, \theta)$ записать в виде

$$r(L, \bar{\theta}, \theta) = E \operatorname{tr} L(\bar{\theta} - \theta)(\bar{\theta} - \theta)^T = \operatorname{tr} LE(\bar{\theta} - \theta)(\bar{\theta} - \theta)^T = \operatorname{tr} LD\bar{\theta}.$$

Ниже приведены три критерия, имеющих ясный статистический смысл и являющихся частными видами L -критерия.

L -оптимальный план с $L = I_m$ называется *A -оптимальным* (от английского термина «average variance» — средняя дисперсия). При A -оптимальном планировании дисперсионная матрица плана имеет наименьший след, а НЛН-оценки неизвестных параметров — минимальную суммарную дисперсию. Благодаря простоте критерий A -оптимальности является одним из наиболее часто используемых.

Прежде чем формулировать следующий критерий, докажем простое вспомогательное утверждение.

Лемма 1.8. Пусть план эксперимента ξ невырожден и имеет вид (1.5), θ — НЛН-оценки параметров θ . Тогда при всех

$x \in X$ $\hat{\theta}^T f(x)$ является НЛН-оценкой функции регрессии $\eta(x) = \theta^T f(x)$, а дисперсия этой оценки равна $\frac{\sigma^2}{N} d(x, \xi)$, где

$$d(x, \xi) = f^T(x) D(\xi) f(x). \quad (1.35)$$

Доказательство. Из теоремы 1.4 гл. 1 следует, что оценка $\hat{\theta}^T f(x)$ является наилучшей в множестве линейных несмещенных оценок величины $\theta^T f(x)$. Подсчитаем дисперсию этой оценки:

$$\begin{aligned} D\hat{\theta}^T f(x) &= E(\hat{\theta}^T f(x) - \theta^T f(x))^2 = E[(\hat{\theta} - \theta)^T f(x)]^2 = \\ &= E\left[\sum_{i=1}^m (\hat{\theta}_i - \theta_i) f_i(x)\right]^2 = E\left[\sum_{i=1}^m f_i(x) (\hat{\theta}_i - \theta_i) (\hat{\theta}_j - \theta_j) f_j(x)\right] = \\ &= E[f^T(x) (\hat{\theta} - \theta) (\hat{\theta} - \theta)^T f(x)] = f^T(x) D\hat{\theta} f(x) = \frac{\sigma^2}{N} f^T(x) D(\xi) f(x). \end{aligned}$$

Лемма доказана.

Функцию (1.35) иногда называют *дисперсией оценки поверхности отклика* (т. е. функции регрессии), хотя, как следует из утверждения леммы 1.8, ее правильнее было бы называть *нормированной дисперсией оценки поверхности отклика*.

Следующий критерий — критерий Q -оптимальности (от английского термина «quadratic mean error» — среднеквадратичная ошибка). Q -оптимальным называется план

$$\xi^* = \arg \min_{\xi} \int_Z d(x, \xi) w(x) dx,$$

где функция $d(x, \xi)$ определяется по формуле (1.35), $w(x)$ — отрицательная функция на множестве Z , задающая относительную важность тех или иных точек $x \in Z$. Множество Z может не совпадать с X . В частности, при $Z \subset X$ имеем задачу интерполяции, а в случае $Z \cap X \neq Z$ — задачу экстраполяции.

Q -оптимальные планы являются L -оптимальными для

$$L = \int_Z f(x) f^T(x) w(x) dx \quad (1.36)$$

(проверить!) и используются в тех случаях, когда исследователя в первую очередь интересует точность оценки функции регрессии во всем множестве X или в некотором его подмножестве.

Оптимальным планом для экстраполяции в точку x_0 (эта точка может не принадлежать множеству X) называется план

$$\xi^* = \arg \min_{\xi} d(x_0, \xi), \quad (1.37)$$

который, как нетрудно проверить, является L -оптимальным с $L = f(x_0) f^T(x_0)$.

Для плана (1.37) дисперсия НЛН-оценки функции регрессии в точке x_0 минимальна.

Приведем также три распространенных критерия, которые относятся к классу минимаксных. Первый из них — критерий *С оптимальности*

$$\Phi[D] = \max_{x \in X} f^T(x) D f(x). \quad (1.38)$$

Соответствующий этому критерию план

$$\xi^* = \arg \min_{\xi} [\max_{x \in X} d(x, \xi)] \quad (1.39)$$

называется *G-оптимальным*.

Величина (1.38) называется *обобщенной дисперсией* (general variance). Выпуклость *G*-критерия следует из теоремы 2.4 приложения 1.

При *G*-оптимальном планировании экспериментатор гарантирует, что во всех точках $x \in X$ дисперсия НЛН-оценок функции регрессии не слишком высока — максимальная по $x \in X$ дисперсия минимизируется.

Как показано в § 2 непрерывный *G*-оптимальный план обладает тем замечательным свойством, что совпадает с непрерывным *D*-оптимальным. Отметим, что это свойство дает еще одну статистическую интерпретацию *D*-оптимальному планированию.

План ξ^* называется *E-оптимальным*, если на нем достигается минимум максимального собственного числа дисперсионной матрицы:

$$\xi^* = \arg \min_{\xi} [\lambda_{\max}(D(\xi))] \quad (1.40)$$

(по-английски eigenvalue — собственное число). При *E*-оптимальном планировании минимизируется длина максимальной оси эллипсоида рассеяния НЛН-оценок параметров. В силу теоремы 1.9 из приложения 1 максимальное собственное число невырожденной матрицы совпадает с минимальным собственным числом матрицы, обратной к ней, и поэтому *E*-оптимальным является также план

$$\xi^* = \arg \max_{\xi} [\lambda_{\min}(M(\xi))].$$

Выпуклость критерия *E*-оптимальности

$$\Phi[D] = \lambda_{\max}(D) \quad (1.41)$$

вытекает из теоремы 2.3 приложения 1.

План

$$\xi^* = \arg \min_{\xi} [\max_{i=1, \dots, m} d_{ii}(\xi)] \quad (1.42)$$

называется *MV-оптимальным* (от английского термина «maximal variance» — максимальная дисперсия); функционал

$$\Phi[D] = \max d_{ii} \quad (1.43)$$

называется *MV-критерием*.

В (1.42) и (1.43) d_{ii} — диагональные элементы матрицы D . Смысл MV -оптимального планирования состоит в нахождении такого плана, при котором максимальная из дисперсий оценок неизвестных параметров будет минимальной. Выпуклость MV -критерия следует из теорем 1.5, 2.4 и формулы (2.8) приложения 1.

1.6. Некоторые множества планов. Для формулировки задачи поиска оптимального плана необходимо задать также множество допустимых планов. Ограничения на множество планов диктуются либо реальными условиями эксперимента, либо выбранным методом оптимизации.

Как указывалось выше, наиболее содержательные математические результаты получаются в том случае, когда множеством допустимых планов является Ξ — множество всех непрерывных планов на X .

В практических постановках задачи оптимального планирования множеством допустимых планов обычно является Ξ_N — множество всех дискретных N -точечных планов, т. е. планов вида (1.5). Оптимальные планы в этом случае могут строиться либо путем прямой минимизации критерия на множестве X^N , либо с помощью методов, использующих специфику задачи и описанных в § 4.

Важным частным случаем множества Ξ_N является Ξ_n — множество так называемых насыщенных планов (при насыщенном планировании число измерений равно числу неизвестных параметров регрессии).

В тех случаях, когда число n точек, в которых могут проводиться измерения, фиксировано, а само количество N измерений может быть выбрано достаточно большим, множеством допустимых планов является $\Xi(n)$ — множество планов вида (1.8) при фиксированном n .

Оптимальные планы на $\Xi(n)$ могут строиться путем минимизации критерия на множестве $X^n \times R^{n-1}$ (т. е. на множестве точек $x_1, \dots, x_n$ и весов $p_1, \dots, p_{n-1}$). Точнее, вместо R^{n-1} в указанном множестве нужно поставить

$$S_{n-1} = \left\{ p_1, \dots, p_{n-1} \mid p_i \geq 0 \quad (i=1, \dots, n-1), \quad \sum_{i=1}^{n-1} p_i \leq 1 \right\}.$$

Кроме того, можно считать, что

$$p_1 \geq p_2 \geq \dots \geq p_{n-1},$$

поскольку информационная матрица плана (1.8) не зависит от того, в каком порядке записаны точки $x_1, \dots, x_n$. Для построения оптимальных планов на $\Xi(n)$ может быть также использован метод, описанный в § 4.

Применяется также постановка задачи поиска оптимальных планов на множестве $\Xi_N \cap \Xi(n)$. Методы § 4 могут быть легко переформулированы и для этого случая.

На практике (особенно в задачах планирования экстремального эксперимента, см. гл. 6) часто используются планы из множеств Ξ , Ξ_N или $\Xi(n)$, обладающие дополнительно свойствами ортогональности и (или) ротатабельности.

Невырожденный план называется *ортогональным*, если его дисперсионная матрица диагональна.

При ортогональном планировании все оценки параметров некоррелированы, а главные оси эллипсоида рассеяния направлены по координатным осям в пространстве параметров. Объем вычислительной работы, необходимый для обработки результатов измерений при ортогональном планировании, относительно мал, и даже при больших m и N (ортогональные планы обычно выбираются в множестве Ξ_N) соответствующие вычисления легко могут быть проведены без помощи ЭВМ. Это основная причина широкой распространенности ортогональных планов в практических расчетах.

План ξ называется *ротатабельным* относительно некоторой точки x_0 , если определяемая по формуле (1.34) функция $d(x, \xi)$ не изменяется при вращении плана относительно x_0 , т. е. если для любых таких $x_1, x_2 \in X$, что $\|x_1 - x_0\| = \|x_2 - x_0\|$, имеет место $d(x_1, \xi) = d(x_2, \xi)$.

Ротатабельность плана ξ относительно точки x_0 означает, что дисперсия НЛН-оценки функции регрессии $\theta^T f(x)$ в точке x одинакова для всех точек, удаленных от x_0 на равное расстояние. Ротатабельные планы часто используются в алгоритмах экстремального планирования, в которых функцией регрессии является градиент целевой функции в некоторой точке, оценка функции регрессии нормируется и по ней строится оценка направления наибольшего возрастания целевой функции из указанной точки, а следствием ротатабельности является равноточность оценки направления (по всем направлениям).

Хотя в принципе приближенно оптимальные планы в множестве ортогональных или ротатабельных планов могут быть построены путем прямой минимизации критерия с учетом соответствующих ограничений, ввиду сложности учета этих ограничений экстремальные задачи оказываются настолько сложными, что указанный способ на практике не используется. Для построения оптимальных ортогональных (ротатабельных) планов обычно используют специальные методы, разработанные для определенных классов регрессионных моделей (см., например, § 1 гл. 4).

Упражнения.

1. Для схем регрессии, рассмотренных в примерах 1.1 и 1.2, вычислите дисперсионную матрицу равномерного плана $\xi_0(dx) = \frac{1}{2} dx \quad (-1 \leq x \leq 1)$.

2. Докажите, что всегда найдется по крайней мере один D_2 -оптимальный план вида (1.8) с числом точек $n \leq s(2m - s + 1)/2$.

3. Покажите, что при A -оптимальном планировании эллипсоид рассеяния (1.30) имеет минимальную сумму квадратов длин осей и наименьшую длину диагонали параллелепипеда, описанного около этого эллипсоида.

4. Сформулируйте критерий Q -оптимальности. Проверьте, что Q -оптимальный план является L -оптимальным для матрицы L , определяемой по формуле (1.36). Докажите аналог теоремы 1.5 об инвариантности непрерывных Q -оптимальных планов относительно невырожденной линейной замены базисных функций и неизвестных параметров.

5. Пусть в схеме регрессионного эксперимента (1.1) измерения не являются равноточными, и пусть дисперсия измерения в точке $x \in X$ равна $\sigma^2/\lambda(x)$, где функция $\lambda(x)$ положительна и известна. Покажите, что с помощью преобразования базисных функций $f_i(x) \rightarrow \sqrt{\lambda(x)} f_i(x)$ такая схема сводится к схеме измерений вида (1.1) с равноточными измерениями.

§ 2. Теоремы эквивалентности

2.1. Общий случай. В данном параграфе предполагается, что множество допустимых планов является Ξ , и тем самым рассматривается только задача непрерывного оптимального планирования.

Ниже приведены две формулировки общей теоремы эквивалентности (для Ψ - и Φ -оптимальных планов). Несмотря на то, что любую из этих теорем можно получить как следствие другой, обе теоремы мы докажем. Первое доказательство очень простое, но использует общее необходимое и достаточное условие оптимальности, второе — чисто формальное.

Сначала рассмотрим задачу Ψ -оптимального планирования (1.26) и предположим, что функционал Ψ выпуклый, а множество информационных матриц $\mathfrak{M}$ компактно. Обозначим через $\Delta(M_1, M_2)$ производную по направлению $M_2 - M_1$ в точке M_1 функции $\Psi[M]$:

$$\Delta(M_1, M_2) = \lim_{\alpha \rightarrow 0+} \frac{\Psi[(1-\alpha)M_1 + \alpha M_2] - \Psi[M_1]}{\alpha}$$

(для выпуклого функционала такая производная всегда существует). Из хорошо известного факта теории экстремальных задач*) (см. [2, 34]) вытекает, что необходимым и достаточным условием оптимальности матрицы M^* является выполнение неравенства

$$\inf_{M \in \mathfrak{M}} \Delta(M^*, M) \geq 0, \quad (2.1)$$

что означает неубывание производной по любому направлению в точке M^* .

Для того чтобы получить из (2.1) конструктивный результат, необходимо воспользоваться спецификой задачи планирования и наложить на Ψ кроме условия выпуклости дополнительное условие дифференцируемости. Будем предполагать, что функция $\Psi[M]$ дифференцируема по элементам матрицы M в окрестности точки M^* (более подробно см. § 3 приложения 1).

) Пусть X — выпуклое множество, f — выпуклый функционал на X , $x^ \in X$. Тогда $x^* = \arg \min_{x \in X} f(x)$, если и только если $\inf_{x \in X} \Delta(x^*, x) \geq 0$, где

$$\Delta(x, y) = \lim_{\alpha \rightarrow 0+} \alpha^{-1} [f((1-\alpha)x + \alpha y) - f(x)].$$

Имеем

$$\Delta(M^*, M) = \frac{\partial \Psi[(1-\alpha)M^* + \alpha M]}{\partial \alpha} \Big|_{\alpha=0+} = \text{tr } \dot{\Psi}[M^*](M - M^*),$$

где

$$\dot{\Psi}[M_0] = \frac{\partial \Psi[M]}{\partial M} \Big|_{M=M_0},$$

а само выражение для производной получено по правилу дифференцирования сложной матричной функции (см. (3.9) из приложения 1).

Далее,

$$\begin{aligned} \inf_M \Delta(M^*, M) &= \inf_{\xi} \int \text{tr } \dot{\Psi}[M^*] f^T(x) \xi(dx) - \\ &- \text{tr } M^* \dot{\Psi}[M^*] = \inf_x \int f^T(x) \dot{\Psi}[M^*] f(x) - \text{tr } M^* \dot{\Psi}[M^*]. \end{aligned} \quad (2.2)$$

По существу, отсюда и вытекает следующая теорема.

Теорема 2.1 (теорема эквивалентности). Если Ψ — выпуклый дифференцируемый функционал и множество информационных матриц $\mathfrak{M}$ компактно, то необходимым и достаточным условием Ψ -оптимальности плана ξ^* является выполнение для всех $x \in X$ неравенства

$$\psi(x, \xi^*) \geq \text{tr } M^* \dot{\Psi}[M^*], \quad (2.3)$$

где

$$\psi(x, \xi) = \int f^T(x) \dot{\Psi}[M(\xi)] f(x). \quad (2.4)$$

Кроме того, план ξ^* сосредоточен в тех точках $x \in X$, в которых в (2.3) достигается равенство.

Первая часть теоремы доказана выше, а вторая вытекает из (2.3) и из следующей леммы, примененной к плану $\xi = \xi^*$.

Лемма 2.1. Для любого плана $\xi \in \Xi$ справедливо тождество

$$\int \psi(x, \xi) \xi(dx) = \text{tr } M(\xi) \dot{\Psi}[M(\xi)], \quad (2.5)$$

где функция $\psi(x, \xi)$ определяется по формуле (2.4).

Доказательство. Используя лемму 1.4 из гл. 1, имеем

$$\begin{aligned} \int \psi(x, \xi) \xi(dx) &= \int \int f^T(x) \dot{\Psi}[M(\xi)] f(x) \xi(dx) = \\ &= \int \text{tr } f^T(x) \dot{\Psi}[M(\xi)] f(x) \xi(dx) = \text{tr} \left[\int f(x) f^T(x) \xi(dx) \right] \dot{\Psi}[M(\xi)] = \\ &= \text{tr } M(\xi) \dot{\Psi}[M(\xi)]. \end{aligned}$$

Лемма доказана.

Докажем еще одно вспомогательное утверждение.

Лемма 2.2. Пусть Φ — дифференцируемый функционал, ξ_0 — невырожденный план, $\alpha \in (0, 1)$, $\xi_\alpha = (1-\alpha)\xi_0 + \alpha\xi_1$, где ξ_1 — произвольный план.

Тогда

$$\frac{\partial \Phi[D(\xi_\alpha)]}{\partial \alpha} \Big|_{\alpha=0+} = \text{tr} \left(\bar{\Phi}(\xi_0) D(\xi_0) - \int \varphi(x, \xi_0) \xi_1(dx) \right), \quad (2.6)$$

где

$$\bar{\Phi}(\xi) = \frac{\partial \Phi(D)}{\partial D} \Big|_{D=D(\xi)}, \quad \varphi(x, \xi) = f^T(x) D(\xi) \bar{\Phi}(\xi) D(\xi) f(x). \quad (2.7)$$

Доказательство. По правилу дифференцирования сложной функции (см. (3.9) из приложения 1) получаем

$$\frac{\partial \Phi[D(\xi_\alpha)]}{\partial \alpha} = \text{tr} \left\{ \bar{\Phi}(\xi_\alpha) \frac{\partial D(\xi_\alpha)}{\partial \alpha} \right\}.$$

Дифференцируя по α обе части соотношения

$$M^{-1}(\xi_\alpha) M(\xi_\alpha) = I_m,$$

имеем

$$\frac{\partial D(\xi_\alpha)}{\partial \alpha} M(\xi_\alpha) + D(\xi_\alpha) \frac{\partial M(\xi_\alpha)}{\partial \alpha} = 0,$$

$$\frac{\partial D(\xi_\alpha)}{\partial \alpha} = -D(\xi_\alpha) \frac{\partial M(\xi_\alpha)}{\partial \alpha} D(\xi_\alpha).$$

Формула (2.6) получается теперь из

$$\frac{\partial M(\xi_\alpha)}{\partial \alpha} = \frac{\partial [(1-\alpha) M(\xi_0) + \alpha M(\xi_1)]}{\partial \alpha} = M(\xi_1) - M(\xi_0),$$

$$\begin{aligned} \text{tr} \bar{\Phi}(\xi_0) D(\xi_0) M(\xi_1) D(\xi_0) &= \int \text{tr} \left(\bar{\Phi}(\xi_0) D(\xi_0) f(x) f^T(x) D(\xi_0) \xi_1(dx) \right) = \\ &= \int f^T(x) D(\xi_0) \bar{\Phi}(\xi_0) D(\xi_0) f(x) \xi_1(dx) = \int \varphi(x, \xi_0) \xi_1(dx). \end{aligned}$$

Лемма доказана.

Отметим, что в частном случае $\xi_1 = \xi_x = \left\{ \begin{smallmatrix} x \\ 1 \end{smallmatrix} \right\}$ (т. е. если план ξ_1 сосредоточен в точке x) формула (2.6) принимает вид

$$\frac{\partial \Phi[D(1-\alpha)\xi_0 + \alpha\xi_x]}{\partial \alpha} \Big|_{\alpha=0+} = \text{tr} \bar{\Phi}(\xi_0) D(\xi_0) - \varphi(x, \xi_0). \quad (2.8)$$

Лемма 2.3 (аналог леммы 2.1). Для любого невырожденного плана ξ имеет место равенство

$$\int_X \varphi(x, \xi) \xi(dx) = \text{tr} \bar{\Phi}(\xi) D(\xi), \quad (2.9)$$

где функция $\varphi(x, \xi)$ определяется по формуле (2.7).

Доказательство. Аналогично доказательству леммы 2.1 имеем

$$\begin{aligned} \int \varphi(x, \xi) \xi(dx) &= \int \text{tr} \bar{\Phi}(\xi) D(\xi) f(x) f^T(x) D(\xi) \xi(dx) = \\ &= \text{tr} \bar{\Phi}(\xi) D(\xi) \int f(x) f^T(x) \xi(dx) D(\xi) = \text{tr} \bar{\Phi}(\xi) D(\xi). \end{aligned}$$

Лемма доказана.

Теперь сформулируем и докажем теорему эквивалентности для Φ -оптимальных планов

Теорема 2.2. Пусть Φ — выпуклый дифференцируемый функционал и существует невырожденный Φ -оптимальный план ξ^* . Тогда 1) Φ -оптимальность плана ξ^* (т. е. выполнение (1.21)) эквивалентна тому, что

$$\max_{x \in X} \varphi(x, \xi^*) = \operatorname{tr} D(\xi^*) \Phi(\xi^*), \quad (2.10)$$

где функция $\varphi(x, \xi)$ определяется по формуле (2.7); 2) план ξ^* сосредоточен в тех точках $x \in X$, в которых

$$\varphi(x, \xi^*) = \max_{x \in X} \varphi(x, \xi^*). \quad (2.11)$$

Доказательство Положим $\xi_\alpha = (1 - \alpha)\xi^* + \alpha\xi_1$, где $\alpha \in (0, 1)$, $\xi_1 = \begin{Bmatrix} x \\ 1 \end{Bmatrix}$, и допустим, что план ξ^* оптимален. Тогда для всех $\alpha \in (0, 1)$ выполняется неравенство

$$\Phi[D(\xi_\alpha)] \geq \Phi[D(\xi^*)].$$

Поэтому

$$\left. \frac{\partial \Phi[D(\xi_\alpha)]}{\partial \alpha} \right|_{\alpha=0+} \geq 0.$$

В силу (2.8) для $\xi_0 = \xi^*$ это неравенство переписывается в виде

$$\varphi(x, \xi^*) \leq \operatorname{tr} \Phi(\xi^*) D(\xi^*). \quad (2.12)$$

Из этого неравенства (используя (2.9) для $\xi = \xi^*$) получаем соотношение (2.10) и второе утверждение теоремы.

Пусть теперь выполняется (2.10). Покажем, что план ξ^* оптимальный. Допустим, что план ξ^* неоптимален, и пусть оптимальным является план ξ_1 . Положим

$$\xi_\alpha = (1 - \alpha)\xi^* + \alpha\xi_1.$$

Из выпуклости Φ и определения ξ_1 имеем

$$\begin{aligned} \Phi[D(\xi_\alpha)] &\leq (1 - \alpha)\Phi[D(\xi^*)] + \alpha\Phi[D(\xi_1)], \\ \left. \frac{\partial \Phi[D(\xi_\alpha)]}{\partial \alpha} \right|_{\alpha=0+} &\leq \Phi[D(\xi_1)] - \Phi[D(\xi^*)] < 0. \end{aligned}$$

Используя (2.6), отсюда получаем

$$\left. \frac{\partial \Phi[D(\xi_\alpha)]}{\partial \alpha} \right|_{\alpha=0+} = \operatorname{tr} \Phi(\xi^*) D(\xi^*) - \int \varphi(x, \xi^*) \xi_1(dx) < 0.$$

С другой стороны, из (2.10) вытекает справедливость неравенства (2.12) для всех $x \in X$. Интегрируя (2.12) по мере $\xi_1(dx)$, получаем

$$\int \varphi(x, \xi^*) \xi_1(dx) \leq \operatorname{tr} \Phi(\xi^*) D(\xi^*), \quad (2.13)$$

что противоречит предыдущему неравенству. Теорема доказана.

2.2. Теорема эквивалентности Кифера — Вольфовица. Задачу поиска D -оптимального плана (1.27) представим в виде задачи Ψ -оптимального планирования с

$$\Psi[M] = -\ln \det M.$$

Из формулы (3.19) приложения 1 получаем

$$\frac{\partial (-\ln \det M)}{\partial M} = -M^{-1},$$

откуда для любого плана ξ

$$\operatorname{tr} M(\xi) \overset{\circ}{\Psi}(\xi) = -\operatorname{tr} I_m = -m,$$

$$\psi(x, \xi) = -f^T(x) M^{-1}(\xi) f(x) = -d(x, \xi)$$

(функция $d(x, \xi)$ та же, что и в § 1). Поэтому из теоремы 2.1 с учетом теоремы 1.4 вытекает следующий классический результат.

Теорема 2.3 (Кифера — Вольфовица). *Если множество информационных матриц Ξ компактно, то следующие утверждения эквивалентны:*

$$a) \xi^* = \arg \max_{\xi \in \Xi} \det M(\xi),$$

т. е. план ξ^ D -оптимален;*

$$б) \xi^* = \arg \min_{\xi \in \Xi} \{\max_{x \in X} d(x, \xi)\},$$

т. е. план ξ^ G -оптимален;*

$$в) \max_{x \in X} d(x, \xi^*) = m.$$

Информационные матрицы всех планов, удовлетворяющих одному из трех указанных утверждений, совпадают между собой. В точках x_i этих планов $d(x_i, \xi^) = m$.*

Теорема 2.3 утверждает, в частности, что D - и G -оптимальные планы совпадают, т. е. задачи D - и G оптимального планирования эквивалентны. Поэтому теорема 2.3 называется теоремой эквивалентности. Более общие теоремы 2.1, 2.2, а также аналогичные теоремы для других критериев были доказаны позднее, но и они по аналогии с теоремой 2.3 были названы теоремами эквивалентности (другое возможное название теорем 2.1, 2.2 — теоремы оптимальности).

Ввиду принципиального значения теоремы 2.3 докажем ее, не обращаясь к приложению 1 и не используя общие теоремы эквивалентности. Сначала докажем лемму, которая будет выполнять роль леммы 2.2.

Лемма 2.4 Пусть $\xi = (1 - \alpha)\xi_0 + \alpha\xi_1$, где $\alpha \in (0, 1)$, план ξ_0 невырожден, а ξ_1 произволен. Тогда

$$\left. \frac{\partial \ln \det M(\xi_\alpha)}{\partial \alpha} \right|_{\alpha=0} = \int_X d(x, \xi_0) \xi_1(dx) - m. \quad (2.14)$$

Доказательство. Ниже используются следующие два хорошо известных соотношения: о разложении определителя невырожденной матрицы $M = \|m_{ij}\|_{i,j=1}^m$ по элементам строки i и о представлении элементов обратной от нее матрицы $M^{-1} = \|m^{ij}\|_{i,j=1}^m$

$$\det M = \sum_{j=1}^m (-1)^{i+j} |M_{ij}| m_{ij}, \quad (2.15)$$

$$m^{ij} = (-1)^{i+j} |M_{ji}| / \det M, \quad i, j = 1, \dots, m. \quad (2.16)$$

Здесь $|M_{ij}|$ — определитель матрицы, полученной из матрицы M вычеркиванием i -й строки и j -го столбца.

По теореме о дифференцировании сложной функции имеем

$$\begin{aligned} \frac{\partial \ln \det M(\xi_a)}{\partial a} &= \frac{1}{\det M(\xi_a)} \frac{\partial \det M(\xi_a)}{\partial a} = \\ &= \frac{1}{\det M(\xi_a)} \sum_{i,j=1}^m \frac{\partial \det M(\xi_a)}{\partial m_{ij}} \frac{\partial m_{ij}(\xi_a)}{\partial a}, \end{aligned}$$

где $m_{ij} = m_{ij}(\xi_a)$ — элементы матрицы $M(\xi_a)$. Из соотношения (2.15) получаем

$$\frac{\partial \det M}{\partial m_{ij}} = (-1)^{i+j} |M_{ji}|$$

Используя (2.16), получаем

$$\frac{1}{\det M} \frac{\partial \det M}{\partial m_{ij}} = m^{ij}.$$

Кроме того,

$$\frac{\partial m_{ij}(\xi_a)}{\partial a} = \frac{\partial [(1-a)m_{ij}(\xi_0) + am_{ij}(\xi_1)]}{\partial a} = m_{ij}(\xi_1) - m_{ij}(\xi_0).$$

Следовательно, используя симметричность информационной матрицы, имеем

$$\begin{aligned} \frac{\partial \ln \det M(\xi_a)}{\partial a} \Big|_{a=0} &= \sum_{i,j=1}^m m^{ij}(\xi_0) [m_{ij}(\xi_1) - m_{ij}(\xi_0)] = \\ &= \sum_{i,j=1}^m m^{ij}(\xi_0) [m_{ij}(\xi_1) - m_{ij}(\xi_0)] = \operatorname{tr} M^{-1}(\xi_0) [M(\xi_1) - M(\xi_0)] = \\ &= \operatorname{tr} M^{-1}(\xi_0) M(\xi_1) - \operatorname{tr} I_m = \operatorname{tr} M^{-1}(\xi_0) \int f(x) f^T(x) \xi_1(dx) - m = \\ &= \int f^T(x) M^{-1}(\xi_0) f(x) \xi_1(dx) - m = \int d(x, \xi_0) \xi_1(dx) - m. \end{aligned}$$

Лемма доказана.

Следующее утверждение, по сути, является частным случаем леммы 2.1.

Лемма 2.5. Для любого невырожденного плана ξ^* справедливо

$$\int d(x, \xi^*) \xi^*(dx) = m, \quad (2.17)$$

$$\max_{x \in X} d(x, \xi^*) \geq m. \quad (2.18)$$

Доказательство. Аналогично доказательству леммы 2.1 имеем

$$\begin{aligned} \int d(x, \xi^*) \xi^*(dx) &= \int f^T(x) D(\xi^*) f(x) \xi^*(dx) = \\ &= \text{tr } D(\xi^*) \int f(x) f^T(x) \xi^*(dx) = \text{tr } D(\xi^*) M(\xi^*) = \text{tr } I_m = m, \end{aligned}$$

т. е. (2.17) выполняется. Неравенство (2.18) следует из (2.17). Лемма доказана.

Доказательство теоремы 2.3. Обозначим через ξ_α план $\xi_\alpha = (1 - \alpha)\xi^* + \alpha\xi_1$, где $\alpha \in (0, 1)$. Сначала предположим, что выполнено а), т. е. план ξ^* D -оптимален. Тогда $\det M(\xi^*) \geq \det M(\xi_\alpha)$, и поэтому

$$\left. \frac{\partial \ln \det M(\xi_\alpha)}{\partial \alpha} \right|_{\alpha=0+} \leq 0.$$

Применяя (2.14) для $\xi_1 = \xi^*$, $\xi_1 = \xi_x = \left\{ \begin{smallmatrix} x \\ 1 \end{smallmatrix} \right\}$, для всех $x \in X$ имеем

$$\left. \frac{\partial \ln \det M(\xi_\alpha)}{\partial \alpha} \right|_{\alpha=0} = d(x, \xi^*) - m \leq 0.$$

Используя (2.18), отсюда получаем в) и G -оптимальность плана ξ^* . Следовательно, доказали, что из а) следуют б), в).

Пусть теперь выполнено б), и предположим, что G -оптимальный план ξ^* , для которого

$$\max_{x \in X} d(x, \xi^*) \leq m \quad (2.19)$$

(для D -оптимального плана в (2.19) будет равенство, следовательно, для G -оптимального — неравенство), не является D -оптимальным (т. е. не выполнено а)). Тогда в качестве ξ_1 выберем D -оптимальный план. Из выпуклости функционала $(-\ln \det M(\xi))$ имеем

$$\ln \det M(\xi_\alpha) \geq (1 - \alpha) \ln \det M(\xi^*) + \alpha \ln \det M(\xi_1).$$

Поэтому

$$\left. \frac{\partial \ln \det M(\xi_\alpha)}{\partial \alpha} \right|_{\alpha=0+} \geq \ln \det M(\xi_1) - \ln \det M(\xi^*) > 0.$$

Применяя (2.14) для $\xi = \xi^*$, получаем строгое неравенство

$$\int d(x, \xi^*) \xi_1(dx) > m,$$

которое противоречит (2.19). Это означает, что из б) следует а), а поэтому и в). Эквивалентность утверждений а), б), в) доказана.

То, что в точках x_i D -оптимального плана ξ^* имеет место $d(x_i, \xi^*) = m$, следует из утверждения в) и соотношения (2.17). Тот факт, что информационные матрицы всех D -оптимальных планов совпадают, вытекает из строгой выпуклости критерия D -оптимальности ($-\ln \det M$) и теорем 1.2, 1.3. Теорема доказана.

2.3. Обобщенный D -критерий. Для обобщенного D -критерия (1.32) функция (2.7) имеет вид

$$\varphi(x, \xi) = f^T(x) D(\xi) A [A^T D(\xi) A]^{-1} A^T D(\xi) f(x),$$

а правая часть (2.10) в случае, когда план ξ^* невырожден, равна

$$\text{tr } D(\xi^*) \overset{\circ}{\Phi}(\xi^*) = s.$$

Поэтому, если множество $\mathfrak{M}$ компактно и план ξ^* невырожден, то, согласно теореме 2.2, утверждения

$$\text{а) } \xi^* = \arg \min_{\xi \in \Xi} \det [A^T D(\xi) A];$$

$$\text{б) } \xi^* = \arg \min_{\xi \in \Xi} [\max_{x \in X} \varphi(x, \xi)];$$

$$\text{в) } \max_{x \in X} \varphi(x, \xi^*) = s$$

эквивалентны.

Это и есть теорема эквивалентности для критерия обобщенной D -оптимальности.

2.4. Линейные критерии. Для L -критерия $\text{tr } LD$ функция (2.7) имеет вид

$$\varphi(x, \xi) = f^T(x) D(\xi) L D(\xi) f(x), \quad (2.20)$$

а правая часть (2.5) (для случая невырожденного плана ξ^*) — вид

$$\text{tr } D(\xi^*) \overset{\circ}{\Phi}(\xi^*) = \text{tr } L D(\xi^*). \quad (2.21)$$

Формулы (2.20) и (2.21) следуют из того, что для $L \geq 0$

$$\frac{\partial \text{tr } LD}{\partial D} = L. \quad (2.22)$$

Эту формулу легко проверить непосредственно (проверьте!), но можно сослаться и на формулу (3.10) из приложения 1.

Таким образом, теорема эквивалентности для L -оптимальных планов имеет следующий вид.

Теорема 2.4. Пусть множество $\mathfrak{M}$ компактно и существует невырожденный L -оптимальный план ξ^* (если $\text{rang } L = m$, то

оптимальный план всегда невырожден). Тогда следующие утверждения эквивалентны:

$$a) \xi^* = \arg \max_{\xi \in S} \text{tr } LD(\xi);$$

$$б) \max_{x \in X} \varphi(x, \xi^*) = \text{tr } LD(\xi^*)$$

(здесь $\varphi(x, \xi)$ определяется по (2.20)). При этом в точках L -оптимального плана ξ^* имеет место равенство

$$\varphi(x_i, \xi^*) = \text{tr } LD(\xi^*).$$

Для критерия A -оптимальности функция (2.20) имеет вид

$$\varphi(x, \xi) = f^T(x) D^2(\xi) f(x).$$

Для Q -оптимальности:

$$\varphi(x, \xi) = \int_Z d^2(x, z, \xi) w(z) dz,$$

где функция

$$d(x, z, \xi) = f^T(x) D(\xi) f(z)$$

пропорциональна ковариации НЛН-оценок функции регрессии в точках x, z .

Для экстраполяции в точку:

$$\varphi(x, \xi) = d^2(x, x_0, \xi).$$

2.5. Теорема эквивалентности для минимаксных критериев. Критерий оптимальности $\Psi[M]$ будем называть минимаксным, если

$$\Psi^*[M] = \max_{u \in U} \Psi[M, u], \quad (2.23)$$

где множество U — компакт, функция $\Psi[M, u]$ непрерывна по u , и при всех $u \in U$ Ψ является выпуклым дифференцируемым по элементам матрицы M функционалом на множестве информационных матриц $\mathfrak{M}$, которое будем предполагать компактным. Минимаксными являются, в частности, рассмотренные в § 1 критерии E -, G - и MV -оптимальности.

Согласно хорошо известному результату выпуклого анализа [34], производная по направлению $M - M^*$ из точки M^* для минимаксного критерия (2.23) равна

$$\Delta(M^*, M) = \sup_{u \in U(\xi^*)} \frac{\partial}{\partial \alpha} \Psi[(1 - \alpha) M^* + \alpha M, u], \quad (2.24)$$

где $U(\xi)$ — подмножество множества U , состоящее из точек

$$u^* = \arg \max_{u \in U} \Psi[M(\xi), u].$$

т. е. решений экстремальной задачи $\max_{u \in U} \Psi[M(\xi), u]$

С учетом (2.24) и полученного в начале параграфа выражения для производной по направлению дифференцируемого критерия необходимое и достаточное условие оптимальности (2.1) записывается в виде

$$\sup_{\xi} \inf_{x \in X} \int_{U(\xi^*)} [\Psi(x, \xi^*, u) - \text{tr } M \Psi^*(M, u) |_{M=M(\xi^*)}] \zeta(du) > 0, \quad (2.25)$$

где супремум берется по множеству вероятностных мер, сосредоточенных на $U(\xi)$.

$$\psi(x, \xi, u) = f^T(x) \Psi(M(\xi), u) / (x),$$

$$\frac{\partial}{\partial M} \Psi(M, u) = \frac{\partial \Psi(M, u)}{\partial M}$$

Упражнения.

1. Сформулируйте теорему эквивалентности для критерия Q -оптимальности.

2. Покажите, что непрерывный план ξ^* является одновременно D - и A -оптимальным, если для некоторой константы $c > 0$ выполнено $cD(\xi^*) = -D^2(\xi^*)$. Обобщите это утверждение для произвольного критерия L -оптимальности ($L \neq I_m$).

3. Сформулируйте теорему эквивалентности для критериев E - и MV -оптимальности.

§ 3. Некоторые следствия теорем эквивалентности

3.1. Оптимальные планы для полиномиальной регрессии на отрезке. Сначала рассмотрим задачу построения непрерывных D -оптимальных планов для полиномиальной регрессии на отрезке, т. е. для схемы измерений (1.1) с $X = [-1, 1]$:

$$f_i(x) = x^{i-1}, \quad i = 1, \dots, m. \quad (3.1)$$

Функция $d(x, \xi)$ имеет вид

$$d(x, \xi) = f^T(x) D(\xi) f(x) = \sum_{i,j=1}^m d_{ij}(\xi) x^{i+j-1}, \quad (3.2)$$

где $d_{ij}(\xi)$ ($i, j = 1, \dots, m$) — элементы матрицы $D(\xi)$. При фиксированном плане ξ функция (3.2) представляет собой многочлен степени $2m-2$, ее производная — полином степени $2m-3$, поэтому $d(x, \xi)$ как функция от x имеет на интервале $(-1, 1)$ не более чем $2m-3$ локальных экстремумов. Учитывая возможные локальные экстремумы в точках 1 и -1 , получаем, что общее количество локальных экстремумов функции (3.2) не превосходит $2m-1$.

Поскольку D -оптимальный план всегда невырожденный и в силу теоремы 2.3 сосредоточен в точках локальных максимумов функции (3.2), а локальные максимумы функции (3.2) чередуются с локальными минимумами, то функция (3.2) (при D -оптимальном плане ξ) на промежутке $[-1, 1]$ имеет ровно m локальных максимумов, два из которых — в точках 1 , -1 . Таким образом, D -оптимальный план имеет вид (1.8) с $n = m$, $-1 = x_1 < x_2 < \dots < x_{m-1} < x_m = 1$.

Из формулы (1.17) приложения 1 следует, что для любого плана указанного вида

$$\det M(\xi) = \prod_{i=1}^m p_i (\det F)^2, \quad (3.3)$$

где

$$F = \|f_i(x_j)\|_{i,j=1}^m.$$

Из (3.3) сразу следует, что для D -оптимального плана

$$p_1^* = p_2^* = \dots = p_m^* = 1/m. \quad (3.4)$$

Далее

$$\det F = \det \|x_i^{j-1}\|_{i,j=1}^m = \prod_{\substack{i,j=1 \\ i \neq j}}^m (x_i - x_j). \quad (3.5)$$

Продифференцируем (3.5) по x_i ($i = 1, \dots, m$) и производные приравняем нулю. Получим, что точки D -оптимального плана являются решением системы уравнений

$$\frac{1}{x_i - x_1} + \dots + \frac{1}{x_i - x_{i-1}} + \frac{1}{x_i - x_{i+1}} + \dots + \frac{1}{x_i - x_m} = 0. \quad (3.6)$$

Положим

$$\varphi(x) = (x - x_2) \dots (x - x_{m-1}).$$

Тогда (3.6) примет вид

$$\frac{1}{x_i + 1} + \frac{1}{x_i - 1} + \frac{\varphi''(x_i)}{\varphi'(x_i)} = 0, \quad i = 2, \dots, m-1, \quad (3.7)$$

или

$$(x_i^2 - 1) \varphi''(x_i) + 2x_i \varphi'(x_i) = 0, \quad i = 2, \dots, m-1$$

Последнее равенство говорит о том, что многочлен

$$(x^2 - 1) \varphi''(x) + 2x \varphi'(x)$$

обращается в нуль в нулях многочлена $\varphi(x)$ и имеет такую же степень $m-2$. Следовательно,

$$(x^2 - 1) \varphi''(x) + 2x \varphi'(x) = \text{const } \varphi(x).$$

Сравнивая коэффициенты при x^{m-2} , находим, что $\text{const} = 3$ и функция $\varphi(x)$ является решением дифференциального уравнения

$$(x^2 - 1) \varphi''(x) + 2x \varphi'(x) - 3\varphi(x) = 0.$$

Из теории ортогональных многочленов известно (см. [37]), что это уравнение на множестве многочленов имеет единственное с точностью до постоянного множителя решение — производную m -го полинома Лежандра:

$$\varphi(x) = P'_m(x).$$

Следовательно, D -оптимальный план единственен и сосредоточен с равными весами в нулях полинома

$$(1 - x^2) \varphi(x) = (1 - x^2) P'_m(x).$$

Таким образом, доказана следующая теорема.

Теорема 3.1. Непрерывный D -оптимальный план для полиномиальной регрессии на отрезке $[-1, 1]$ единственен и сосредоточен с равными весами в m точках, являющихся нулями полинома $(1 - x^2) P'_m(x)$, где $P_m(x)$ — полином Лежандра.

Отметим, что из доказательства теоремы вытекает способ вычисления верхней границы числа l локальных максимумов функции $d(x, \xi)$ для случая, когда базисные функции $f_i(x)$ ($i = 1, \dots, m$) — произвольные полиномы на отрезке $X = [a, b]$. Пусть K — максимальная степень указанных полиномов, тогда для любого плана ξ степень полинома $d(x, \xi)$ равна $2K$, число локальных экстремумов функции $d(x, \xi)$ (включая экстремумы в точках a, b) не превосходит $2K + 1$ и, следовательно, $l \leq K + 1$. Для случая, рассмотренного в теореме, $K = m - 1$ и $l = m$. Подобные рассуждения нужны при выборе алгоритмов поиска максимума функции $d(x, \xi)$, которые являются составной частью алгоритмов из § 4.

Из утверждения теоремы 3.1 и известных результатов об асимптотическом распределении нулей ортогональных полиномов (см. [37]) вытекает тот факт, что при возрастании m последовательность D -оптимальных планов ξ_m^* для полиномиальной регрессии степени $m - 1$ на отрезке $[-1, 1]$ слабо сходится к плану

$$\xi_\infty^*(dx) = \frac{1}{\pi \sqrt{1-x^2}} dx, \quad (3.8)$$

т. е. к вероятностной мере с плотностью

$$p(x) = \frac{1}{\pi \sqrt{1-x^2}}, \quad x \in [-1, 1].$$

План (3.8) можно использовать в случае, когда точная степень полинома неизвестна, чтобы оценивать параметры сразу нескольких полиномиальных регрессий с целью выбора среди них наиболее подходящей. В этой связи возникает вопрос о том, как ведут себя характеристики точности предельного плана ξ_∞^* при его использовании для полинома точной степени n .

Пусть $d_n(x, \xi) = f_{(n)}^T(x) M_{(n)}^{-1}(\xi) f_{(n)}(x)$ — дисперсия оценки полинома точной степени n в точке x по наблюдениям плана ξ ; $f_{(n)}^T(x) = (1, x, \dots, x^n)$ — вектор базисных функций; $M_{(n)}(\xi) = \int_{-1}^1 f_{(n)}(x) f_{(n)}^T(x) \xi(dx)$ — информационная матрица для регрессии степени n ; $d_n(\xi) = \max_{x \in [-1, 1]} d_n(x, \xi)$ — максимальное значение дисперсии на интервале наблюдения.

Используя инвариантность величины $d_n(x, \xi)$ относительно выбора базиса (см. теорему 1.5), переходим от полиномов $1, x, x^2, \dots$ к ортонормированным относительно меры ξ_∞^* полиномам Чебышева 1 рода:

$$1, \sqrt{2} T_1(x), \dots, \sqrt{2} T_n(x), \\ T_k(\cos \varphi) = \cos k\varphi, \quad k = 1, \dots, n.$$

Получим

$$d_n(x, \xi) = 1 + 2 \sum_{i=1}^n T_i^2(x) = n + \frac{1}{2} + \frac{1}{2} U_{2n}(x),$$

где $U_k(\cos \varphi) = \sin(k+1)\varphi / \sin \varphi$ — полином Чебышева II рода от переменной $x = \cos \varphi$. Отсюда $d_n(\xi_\infty^*) = 2n + 1$, что примерно вдвое больше оптимального значения $d_n(\xi_n^*) = n + 1$.

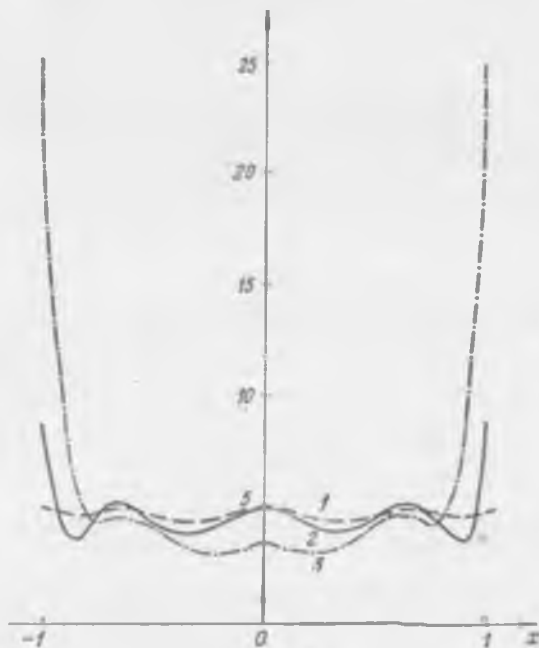


Рис. 7

В качестве иллюстрации на рис. 7 для случая $n = 4$ приведены графики функций $d(x, \xi_n^*)$, $d(x, \xi_\infty^*)$ и $d(x, \xi_0)$ (обозначены соответственно цифрами 1, 2, 3), где ξ_0 — равномерный план

$$\xi_0(dx) = \frac{1}{2} dx, \quad x \in [-1, 1]. \quad (3.9)$$

Обратимся к простейшим видам полиномиальной регрессии на отрезке $X = [-1, 1]$.

Пример 3.1. Дан полином нулевой степени $\eta(x, \theta) = \theta_1$. Очевидно, что для любого плана ξ выполнено равенство $M(\xi) = D(\xi) = 1$; поэтому любой план является оптимальным.

Пример 3.2. Дана регрессия первого порядка

$$\eta(x, \theta) = \theta_1 + \theta_2 x, \quad f(x) = (1, x)^T, \quad X = [-1, 1].$$

Для любого плана ξ

$$M(\xi) = \begin{bmatrix} 1 & \int_{-1}^1 x \xi(dx) \\ \int_{-1}^1 x \xi(dx) & \int_{-1}^1 x^2 \xi(dx) \end{bmatrix},$$

$$\det M(\xi) = \int_{-1}^1 x^2 \xi(dx) - \left[\int_{-1}^1 x \xi(dx) \right]^2. \quad (3.10)$$

Первый член в правой части (3.10) не больше единицы (так как $x \in [-1, 1]$), а второй не больше нуля. Поэтому, если найдется план ξ^* , для которого

$$\int_{-1}^1 x^2 \xi^*(dx) = 1, \quad \int_{-1}^1 x \xi^*(dx) = 0,$$

то этот план и будет оптимальным. Первое из этих условий показывает, что план может быть сосредоточен только в точках $1, -1$, а второе — что этим точкам он может приписывать только равные веса. Такой план существует и имеет вид

$$\xi^* = \left\{ \begin{matrix} -1 & 1 \\ 1/2 & 1/2 \end{matrix} \right\}. \quad (3.11)$$

Проверим теперь D -оптимальность плана (3.11) по теореме эквивалентности Кифера — Вольфовица. Имеем

$$M(\xi^*) = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}, \quad d(x, \xi^*) = (1, x) M^{-1}(\xi^*) (1, x)^T = 1 + x^2.$$

Максимум функции $1 + x^2$ равен 2 и достигается только в точках $1, -1$, и поэтому план (3.11) является D -оптимальным.

Пример 3.3. Дана квадратичная регрессия

$$\eta(x, \theta) = \theta_1 + \theta_2 x + \theta_3 x^3, \quad f(x) = (1, x, x^2)^T, \quad (3.12)$$

на отрезке $[-1, 1]$. Проверим с помощью теоремы 2.3, что план

$$\xi^* = \left\{ \begin{matrix} -1 & 0 & 1 \\ 1/3 & 1/3 & 1/3 \end{matrix} \right\} \quad (3.13)$$

D -оптимален. Действительно,

$$M(\xi^*) = \begin{bmatrix} 1 & 0 & 2/3 \\ 0 & 2/3 & 0 \\ 2/3 & 0 & 2/3 \end{bmatrix}, \quad M^{-1}(\xi^*) = \begin{bmatrix} 3 & 0 & -3 \\ 0 & 3/2 & 0 \\ -3 & 0 & 9/2 \end{bmatrix},$$

$$d(x, \xi^*) = 3 - \frac{9}{2} x^2 (1 - x^2).$$

Максимум функции $3 - \frac{9}{2}x^2(1-x^2)$ равен 3 и достигается только в точках $-1, 0, 1$.

Более сложно построить дискретный D -оптимальный план для случая, когда N не кратно 3. Можно доказать (доказательство здесь не приводится), что при любом N дискретный D -оптимальный план для квадратичной регрессии на отрезке $[-1, 1]$, так же как и непрерывный, всегда сосредоточен в точках $-1, 0, 1$.

В связи с обсуждением дискретных оптимальных планов в следующем примере покажем, что в отличие от непрерывного случая дискретные D -оптимальные планы могут не совпадать с дискретными G -оптимальными планами.

В этом легко убедиться на следующем примере.

Пример 3.4 (продолжение примера 3.2). Пусть $X = [-1, 1]$, $\eta(x) = \theta_1 + \theta_2 x$, и пусть число измерений задано и равно 3. Сначала найдем дискретный D -оптимальный план ξ вида

$$\xi = \left\{ \begin{matrix} x_1 & x_2 & x_3 \\ 1/3 & 1/3 & 1/3 \end{matrix} \right\}. \quad (3.14)$$

где по крайней мере две из точек x_1, x_2, x_3 различны (иначе план (3.14) является вырожденным).

По определению

$$M(\xi) = \sum_{j=1}^3 \frac{1}{3} f(x_j) f^T(x_j) = \frac{1}{3} \begin{bmatrix} 3 & x_1 + x_2 + x_3 \\ x_1 + x_2 + x_3 & x_1^2 + x_2^2 + x_3^2 \end{bmatrix},$$

$$\det M(\xi) = \frac{2}{9} (x_1^2 + x_2^2 + x_3^2 - x_1 x_2 - x_1 x_3 - x_2 x_3).$$

Приравнивая частные производные от $\det M(\xi)$ по x_j ($j = 1, 2, 3$) нулю, получаем систему уравнений

$$\begin{aligned} 2x_1 - x_2 - x_3 &= 0, \\ -x_1 + 2x_2 - x_3 &= 0, \\ -x_1 - x_2 + 2x_3 &= 0, \end{aligned}$$

решением которой является $x_1 = x_2 = x_3 = c$, где c — произвольная константа. Для этих значений x_1, x_2, x_3 план (3.14) вырожден и, следовательно, не является D -оптимальным.

Таким образом, в качестве x_1, x_2, x_3 могут быть выбраны только граничные точки $+1, -1$. Для обоих получающихся планов

$$\xi_1 = \left\{ \begin{matrix} -1 & 1 \\ 1/3 & 2/3 \end{matrix} \right\}, \quad \xi_2 = \left\{ \begin{matrix} -1 & 1 \\ 2/3 & 1/3 \end{matrix} \right\}$$

значения $\det M(\xi_i)$ совпадают:

$$\det M(\xi_1) = \det M(\xi_2) = 8/9.$$

Для плана ξ_1 имеем

$$M(\xi_1) = \begin{bmatrix} 1 & 1/3 \\ 1/3 & 1 \end{bmatrix}, \quad D(\xi_1) = \begin{bmatrix} 9/8 & -3/8 \\ -3/8 & 9/8 \end{bmatrix},$$

$$\begin{aligned} d(x, \xi_1) &= \begin{bmatrix} 1 & x \end{bmatrix} \begin{bmatrix} 9/8 & -3/8 \\ -3/8 & 9/8 \end{bmatrix} \begin{bmatrix} 1 \\ x \end{bmatrix} = \\ &= \left[\frac{9}{8} - \frac{3}{8}x \right] x - \frac{3}{8} + \frac{9}{8}x \left[\frac{1}{x} \right] = \frac{9}{8} - \frac{3}{8}x - \frac{3}{8}x + \frac{9}{8}x^2 = \\ &= \frac{1}{8}(9 - 1 + 1 - 6x + 9x^2) = 1 + \frac{1}{8}(3x - 1)^2, \\ \max_{x \in [-1, 1]} d(x, \xi_1) &= d(-1, \xi_1) = 3. \end{aligned}$$

Аналогично проверяется (проверьте!), что

$$\max_{x \in [-1, 1]} d(x, \xi_2) = d(1, \xi_2) = 3.$$

Теперь рассмотрим план

$$\xi_3 = \left\{ \begin{bmatrix} -1 & 0 & 1 \\ 1/3 & 1/3 & 1/3 \end{bmatrix} \right\}.$$

Для этого плана имеем

$$M(\xi_3) = \frac{1}{3} \begin{bmatrix} 3 & 0 \\ 0 & 2 \end{bmatrix}, \quad D(\xi_3) = \begin{bmatrix} 1 & 0 \\ 0 & 3/2 \end{bmatrix},$$

$$d(x, \xi_3) = 1 + \frac{3}{2}x^2,$$

$$\max_{x \in [-1, 1]} d(x, \xi_3) = 5/2 < 3.$$

Следовательно, дискретные D -оптимальные планы ξ_1 и ξ_2 по критерию G -оптимальности хуже плана ξ_3 и поэтому не являются дискретными G -оптимальными.

Покажем, как для построения D -оптимальных планов может быть использована теорема 3.1. Для этого сначала приведем несколько первых многочленов Лежандра:

$$\begin{aligned} P_0(x) &= 1, & P_1(x) &= x, & P_2(x) &= (3x^2 - 1)/2, \\ P_3(x) &= (5x^3 - 3x)/2, & P_4(x) &= (35x^4 - 30x^2 + 3)/8, \\ P_5(x) &= (63x^5 - 70x^3 + 15x)/8, \\ P_6(x) &= (231x^6 - 315x^4 + 105x^2 - 5)/16. \end{aligned}$$

В общем виде многочлен Лежандра P_n определяется по формуле

$$P_n(x) = \frac{1}{n!2^n} \frac{d^n}{dx^n} (x^2 - 1)^n.$$

Пример 3.5. $X = [-1, 1]$, $\eta(x) = \theta_1 + \theta_2 x + \theta_3 x^2 + \theta_4 x^3$. По теореме 3.1 непрерывный D -оптимальный план сосредоточен в четырех точках $x_1 = -1$, x_2 , x_3 , $x_4 = 1$ с равными весами.

Точки x_2^* , x_3^* являются нулями многочлена $P_3'(x) = \frac{1}{2}(5x^3 - 3x)' = \frac{1}{2}(15x^2 - 3) = \frac{3}{2}(5x^2 - 1)$. Следовательно, $x_2^* = -1/\sqrt{5}$, $x_3^* = 1/\sqrt{5}$.

Перейдем к примерам применения теоремы эквивалентности для линейных критериев оптимальности.

Пример 3.6. A -оптимальный план для квадратичной регрессии $\eta(x) = \theta_1 + \theta_2 x + \theta_3 x^2$ на отрезке $[-1, 1]$.

Из соображений симметрии предположим, что A -оптимальный план имеет вид

$$\xi = \left\{ \begin{array}{ccc} -1 & 0 & 1 \\ p/2 & 1-p & p/2 \end{array} \right\} \quad (3.15)$$

и найдем p ($0 < p \leq 1$) из условия минимума функционала $\text{tr } D(\xi)$. Имеем

$$M(\xi) = \begin{bmatrix} 1 & 0 & p \\ 0 & p & 0 \\ p & 0 & p \end{bmatrix}, \quad D(\xi) = \begin{bmatrix} 1/(1-p) & 0 & -1/(1-p) \\ 0 & 1/p & 0 \\ -1/(1-p) & 0 & 1/[p(1-p)] \end{bmatrix},$$

$$\text{tr } D(\xi) = 2/[p(1-p)].$$

Минимум этого выражения достигается при $p = 1/2$ (проверьте!).

Используя теорему эквивалентности (теорему 2.4 для $L = I_m$), покажем, что полученный план

$$\xi^* = \left\{ \begin{array}{ccc} -1 & 0 & 1 \\ 1/4 & 1/2 & 1/4 \end{array} \right\}$$

A -оптимален. Действительно,

$$D(\xi^*) = \begin{bmatrix} 2 & 0 & -2 \\ 0 & 2 & 0 \\ -2 & 0 & 4 \end{bmatrix}, \quad \text{tr } D(\xi^*) = 8,$$

$$\begin{aligned} \varphi(x, \xi^*) &= [1 \quad x \quad x^2] \begin{bmatrix} 2 & 0 & -2 \\ 0 & 2 & 0 \\ -2 & 0 & 4 \end{bmatrix}^2 \begin{bmatrix} 1 \\ x \\ x^2 \end{bmatrix} = \\ &= [2 - 2x^2 \quad 2x \quad -2 + 4x^2] \begin{bmatrix} 2 - 2x^2 \\ 2 \\ -2 + 4x^2 \end{bmatrix} = \\ &= (2 - 2x^2)^2 + 4x^2 + (2 - 4x^2)^2 = 8 - 20x^2(1 - x^2). \end{aligned}$$

Очевидно, что

$$\max_{x \in X} \varphi(x, \xi^*) = 8 = \text{tr } D(\xi^*),$$

и этот максимум достигается в точках $-1, 0, 1$, откуда и следует A -оптимальность плана ξ^* .

Пример 3.7 (продолжение примера 3.6). Пусть в обозначениях предыдущего примера $\Phi[D] = \text{tr } LD$, где L — диагональная матрица с элементами 1, 2, 1. Из соображений симметрии

опять предположим, что оптимальный план имеет вид (3.15), и выберем оптимальное значение p . Имеем

$$\begin{aligned} \text{tr } LD(\xi) &= \text{tr} \begin{bmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 1/(1-p) & 0 & -1/(1-p) \\ 0 & 1/p & 0 \\ -1/(1-p) & 0 & 1/[p(1-p)] \end{bmatrix} = \\ &= 1/(1-p) + 2/p + 1/[p(1-p)]. \end{aligned}$$

Минимизируя эту функцию по p , получаем $p^* = 3 - \sqrt{6}$ (проверьте!).

Проверим L -оптимальность получившегося плана

$$\xi^* = \left\{ \begin{array}{ccc} -1, & 0, & 1 \\ (3 - \sqrt{6})/2, & \sqrt{6} - 2, & (3 - \sqrt{6})/2 \end{array} \right\}$$

с помощью теоремы 2.4. Имеем

$$\begin{aligned} \text{tr } D(\xi^*) &= 5 + 2\sqrt{6}, \\ \varphi(x, \xi^*) &= \frac{32 + 13\sqrt{6}}{3} x^4 - \frac{32 + 13\sqrt{6}}{3} x^2 + 5 + 2\sqrt{6}, \\ \max_{x \in [-1, 1]} \varphi(x, \xi^*) &= \varphi(0, \xi^*) = \varphi(1, \xi^*) = \varphi(-1, \xi^*) = 5 + 2\sqrt{6}. \end{aligned}$$

Следовательно, план ξ^* L -оптимален.

Приведем теперь пример, показывающий, что информационные матрицы D_1 -оптимальных планов могут различаться.

Пример 3.8. Пусть $X = [-1, 1]$, $\eta(x) = \theta_1 + \theta_2 x$,

$$A = \begin{bmatrix} 1 & a \\ 0 & 0 \end{bmatrix}, \quad 0 < a \leq 1, \quad 0 < b \leq 1,$$

$$\xi_1 = \left\{ \begin{array}{cc} -a & a \\ 1/2 & 1/2 \end{array} \right\}, \quad \xi_2 = \left\{ \begin{array}{cc} -b & 0 \\ 1/3 & 1/3 \end{array} \right\}.$$

С помощью теоремы эквивалентности проверим, что планы ξ_1 и ξ_2 D -оптимальны, т. е. минимизируют $\det AD(\xi)A$ по всем планам ξ . Имеем

$$\begin{aligned} M(\xi_1) &= \begin{bmatrix} 1 & 0 \\ 0 & a^2 \end{bmatrix}, & D(\xi_1) &= \begin{bmatrix} 1 & 0 \\ 0 & a^{-2} \end{bmatrix}, \\ M(\xi_2) &= \begin{bmatrix} 1 & 0 \\ 0 & 2b^2/3 \end{bmatrix}, & D(\xi_2) &= \begin{bmatrix} 1 & 0 \\ 0 & 3b^{-2}/2 \end{bmatrix}. \end{aligned}$$

Для обоих планов $\xi = \xi_1$ и $\xi = \xi_2$

$$\varphi(x, \xi) = [1 \ x] AD(\xi) A \begin{bmatrix} 1 \\ x \end{bmatrix} = [1 \ x] \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} 1 \\ x \end{bmatrix} = 1,$$

$$\max_{x \in X} \varphi(x, \xi) = 1,$$

откуда следует D_1 -оптимальность планов ξ_1 и ξ_2 несмотря на то, что их информационные матрицы различны.

3.2. D -оптимальные планы для тригонометрической регрессии на отрезке.

Теорема 3.2. Пусть $X = [0, 2\pi)$, $m = 2k + 1$,

$$\eta(x, \theta) \Rightarrow \theta_1 + \sum_{j=1}^k [\theta_{2j} \cos jx + \theta_{2j+1} \sin jx]. \quad (3.16)$$

Непрерывным D -оптимальным планом для регрессии (3.16) является любой план

$$\xi_N = \left\{ \begin{matrix} x_1^*, \dots, x_N^* \\ 1/N, \dots, 1/N \end{matrix} \right\}, \quad (3.17)$$

где число N четно,

$$x_i^* = \frac{i-1}{N} 2\pi, \quad i = 1, \dots, N, \quad N \geq 2k + 1$$

(x_i^* — равноотстоящие точки на $[0, 2\pi)$), а также равномерный план

$$\xi(dx) = \frac{1}{2\pi} dx. \quad (3.18)$$

Доказательство. В силу ортогональности базисных функций

$$f(x) = (1, \sin x, \cos x, \dots, \sin kx, \cos kx)^T$$

план (3.18) имеет диагональную информационную матрицу

$$M(\xi) = \begin{bmatrix} 1 & 0 & 0 & \dots & 0 \\ 0 & 1/2 & 0 & \dots & 0 \\ 0 & 0 & 1/2 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \dots & 1/2 \end{bmatrix}. \quad (3.19)$$

Действительно, для любого натурального j имеют место

$$\int_0^{2\pi} \sin jx \, dx = 0, \quad \int_0^{2\pi} \cos jx \, dx = 0, \quad (3.20)$$

$$\frac{1}{2\pi} \int_0^{2\pi} \sin^2 jx \, dx = \frac{1}{2\pi} \int_0^{2\pi} \left(\frac{1 - \cos 2jx}{2} \right) dx = \frac{1}{2},$$

$$\frac{1}{2\pi} \int_0^{2\pi} \cos^2 jx \, dx = \frac{1}{2\pi} \int_0^{2\pi} \frac{1 + \cos 2jx}{2} dx = \frac{1}{2}.$$

Представляя произведения синусов и косинусов от аргументов jx и kx ($k \neq j$) через полусумму (полуразность) синусов и косинусов от аргументов $\frac{j-k}{2}x$, $\frac{j+k}{2}x$, получаем нулевые внедиагональные элементы матрицы $M(\xi)$.

Обратная к матрице (3.19) равна

$$D(\xi) = \begin{bmatrix} 1 & 0 & \dots & 0 \\ 0 & 2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 2 \end{bmatrix}.$$

Следовательно,

$$d(x, \xi) = f^T(x) D(\xi) f(x) =$$

$$= 1 + 2(\sin^2 x + \cos^2 x + \dots + \sin^2 kx + \cos^2 kx) = 1 + 2k = m.$$

В силу теоремы 2.3 получаем, что план (3.18) D -оптимален.

Для того чтобы доказать D -оптимальность плана (3.17), достаточно проверить, что его информационная матрица имеет вид (3.19). В силу указанных выше тригонометрических формул достаточно показать, что имеют место аналогичные формулы (3.20):

$$\sum_{i=1}^N \sin jx_i = 0, \quad \sum_{i=1}^N \cos jx_i = 0.$$

Эти формулы для четного N следуют из того, что прямые, направленные из начала координат под углами $\alpha_i = jx_i = \frac{(i-1)j}{N} \cdot 2\pi$ ($i = 0, \dots, N-1$), пересекают единичную окружность в точках, расположенных симметрично относительно обеих координатных осей. Теорема доказана.

Из теоремы 3.2 следует, в частности, что D -оптимальный план для регрессии вида (3.16) является D -оптимальным и для регрессии того же вида, но с меньшим числом неизвестных параметров.

3.3. Пример D -оптимального плана, сосредоточенного в максимально возможном числе точек. Во всех приведенных в данном параграфе примерах оптимальные планы можно было выбирать сосредоточенными в минимально возможном числе точек, равном m . Такой простой ситуации бывает далеко не всегда. Приведем пример случая, когда не существует оптимального плана, сосредоточенного менее чем в $m(m+1)/2$ точках для произвольного $m > 1$ (при $m = 1$ проблемы нет).

Пусть $X = \{x_{ij} | 1 \leq i \leq m, 1 \leq j \leq m\}$ — конечное множество, состоящее из $m(m+1)/2$ точек x_{ij} . Выберем число α такое, что $2(2m-4)/(2m-3) < \alpha < 2$ (например, $\alpha = (4m-7)/(2m-3)$). Положим $f_i(x_{ij}) = \sqrt{\alpha} \delta_{ij}$, $f_i(x_{ij}) = \delta_{ii} + \delta_{ij}$, $i \neq j$, где $i, j, l = 1, \dots, m$, а

$$\delta_{ij} = \begin{cases} 0, & i \neq j, \\ 1, & i = j, \end{cases}$$

есть символ Кронекера. Пусть ξ — произвольный непрерывный план на множестве X , согласно которому вес точки x_{ij} равен p_{ij} ($\sum_{i,j} p_{ij} = 1$, $p_{ij} \geq 0$).

Элементы $m_{ii}(\xi)$ информационной матрицы $M(\xi)$ плана ξ равны

$$m_{ii}(\xi) = p_{ii}\alpha + \sum_{l=1}^{i-1} p_{il} + \sum_{j=i+1}^m p_{ij}, \quad m_{ij}(\xi) = m_{ji}(\xi) = p_{ij}, \quad i < j.$$

Пусть ξ^* есть D -оптимальный план, который точкам x_{ij} приписывает веса p_{ij} . Покажем, что все эти веса p_{ij} строго положительны.

Поскольку критерий D -оптимальности является строго выпуклым функционалом, то информационные матрицы всех D -оптимальных планов совпадают. Из приведенных выражений для элементов информационной матрицы следует, что оптимальный план ξ^* через матрицу $M(\xi^*)$ определяется однозначно и поэтому также единственен.

Далее, D -оптимальные планы инвариантны относительно перестановок параметров (частный случай линейного преобразования параметров). Поэтому

будем рассматривать только класс планов, инвариантных в следующем смысле. Для любой перестановки $\pi = (\pi_1, \dots, \pi_m)$ элементов $1, \dots, m$ определим $\bar{g}\theta_i = \theta_{\pi_i}$,

$$gx_{ij} = \begin{cases} x_{\pi_i \pi_j}, & \text{если } \pi_i \leq \pi_j, \\ x_{\pi_j \pi_i}, & \text{если } \pi_i \geq \pi_j. \end{cases}$$

Тогда

$$(\bar{g}\theta)^T / (gx) = \theta^T / x$$

и инвариантность планов — это инвариантность относительно всевозможных преобразований g указанного вида. Для всех инвариантных планов

$$p_{ii} = a, \quad p_{ij} = b \quad i < j,$$

где параметры a и b удовлетворяют соотношению

$$ma + m(m-1)b/2 = 1. \quad (3.21)$$

Для таких планов

$$M(\xi) = (aa + (m-2)b)I_m + bJ_m. \quad (3.22)$$

J_m — матрица порядка $m \times m$, состоящая из единиц. Определитель матрицы (3.22) равен

$$\det M(\xi) = [aa + 2(m-1)b][aa + (m-2)b]^{m-1}. \quad (3.23)$$

Из (3.21) выразим a и b :

$$a = [1 - m(m-1)b/2]/m, \quad (3.24)$$

$$b = 2(1 - ma)/[m(m-1)]. \quad (3.25)$$

Подставив (3.24) в (3.23) и взяв производную от $\det M(\xi)$ по b , получаем, что эта производная при $b = 0$ отрицательна (в случае $\alpha < 2$). Подставив (3.25) в (3.23) и взяв производную от $\det M(\xi)$ по a , получаем строгую положительность этой производной при $a = 0$ (в случае $\alpha > 2(2m-4)/(2m-3)$). Следовательно, для D -оптимального плана $a \neq 0$, $b \neq 0$. Поэтому D -оптимальный план сосредоточен во всех точках множества X .

Упражнения.

1. Для линейной регрессии $\eta(x) = \theta_0 + \theta_1 x$ на отрезке $[-1, 1]$ сравните по критериям A -, Q -, G -, MV - и E -оптимальности планы

$$\left\{ \begin{matrix} -1 & 1 \\ 1/2 & 1/2 \end{matrix} \right\}, \quad \left\{ \begin{matrix} -1 & 0 & 1 \\ 1/3 & 1/3 & 1/3 \end{matrix} \right\}, \quad \frac{1}{2} dx.$$

Укажите, какие из них являются A - и G -оптимальными.

2. Для квадратичной регрессии $\eta(x) = \theta_0 + \theta_1 x + \theta_2 x^2$ на отрезке $[-1, 1]$:

- постройте дискретный D -оптимальный план для $N = 12$ измерений;
- сравните по критериям A -, Q -, G -, MV - и E -оптимальности планы

$$\left\{ \begin{matrix} -1 & 0 & 1 \\ 1/4 & 1/2 & 1/4 \end{matrix} \right\}, \quad \left\{ \begin{matrix} -1 & 0 & 1 \\ 1/3 & 1/3 & 1/3 \end{matrix} \right\}, \quad \frac{1}{2} dx.$$

в) с помощью теорем эквивалентности проверьте, являются ли выписанные планы оптимальными по указанным критериям.

3. Используя теорему 3.1, постройте D -оптимальные планы для полиномиальной регрессии четвертой степени на отрезке $[-1, 0]$.

4. Не используя теорему 3.1, докажите, что план из примера 3.5 является D -оптимальным.

5. Для квадратичной регрессии $\eta(x) = \theta_0 + \theta_1 x + \theta_2 x^2$ на отрезке $[-1, 1]$ постройте в множестве планов вида (3.15) Q -оптимальный план для $\omega(x) = 1$.

$Z = X$, после чего с помощью теоремы эквивалентности убедиться, что построенный план является Q -оптимальным в множестве Ξ .

6. С помощью теоремы эквивалентности проверьте, что для функции регрессии $\eta(x) = \theta_1 + \theta_2 x + \theta_3 x^2 + \theta_4 x^3$ на отрезке $[-1, 1]$ Q -оптимальным планом при $w(x) = 1$, $Z = X$ является план

$$\xi^* = \begin{Bmatrix} -1 & -1/\sqrt{5} & 1/\sqrt{5} & 1 \\ p & 1/2 - p & 1/2 - p & p \end{Bmatrix}.$$

где $p = 1/(2 + 2\sqrt{5})$.

7. С помощью теоремы эквивалентности проверьте, что для функции регрессии $\eta(x) = \theta_1 + \theta_2 x + \theta_3 x^2$ на отрезке $[-1, 1]$ оптимальным планом для экстраполяции в точку $x_0 = 2$ является план

$$\xi^* = \begin{Bmatrix} -1 & 0 & 1 \\ 1/7 & 3/7 & 3/7 \end{Bmatrix}.$$

Сравните этот план с D -оптимальным.

8. Постройте D -, G - и A -оптимальные планы для функции регрессии $\eta(x) = \theta_1 + \theta_2 \sin x$ на отрезке $[0, 2\pi]$.

9. Для регрессии $\eta(x) = \theta_1 + \theta_2 \sin x + \theta_3 \cos x$ на отрезке $[0, 2\pi]$ постройте дискретный D -оптимальный план для проведения $N = 5$ измерений. Проверьте, является ли построенный план непрерывным A -оптимальным.

10. Используя формулы тригонометрии, покажите, что требование четности числа N в теореме 3.2 излишне.

11. Покажите, что планы (3.17) и (3.18) являются A - и Q -оптимальными (для $w(x) = 1$, $Z = X$) по оцениванию параметров функции регрессии (3.16).

12. Пусть множество планирования

$$X = \left\{ x = (x(1), \dots, x(k))^T \mid \sum_{i=1}^k [x(i)]^2 \leq 1 \right\}$$

есть единичный k -мерный шар, а функция регрессии

$$\eta(x) = \theta_1 + \sum_{i=1}^k \theta_{i+1} x(i) + \sum_{i < j} \theta_{ij} x(i) x(j)$$

есть полином второго порядка (здесь число параметров равно $m = (k+1)(k+2)/2$). Доказать (сначала для случаев $k=2$ и $k=3$), что D -оптимальным является план $\xi^* = \alpha \xi_1 + (1-\alpha)\xi_2$, где $\alpha = 1/m$, план ξ_1 сосредоточен в нуле, а план ξ_2 представляет собой равномерную меру на поверхности сферы

$$\left\{ x = (x(1), \dots, x(k))^T \mid \sum_{i=1}^k [x(i)]^2 = 1 \right\}.$$

§ 4. Численные методы построения оптимальных планов

4.1. Построение непрерывных оптимальных планов. Рассмотрим некоторые особенности исходной экстремальной задачи поиска минимума $\Psi[M(\xi)]$. Эта задача, как ясно из следствия 1.1, может быть сведена к конечномерной задаче с размерностью не более чем $(1/2)m(m+1)(k+1)$, где m — число неизвестных параметров, а k — размерность X ($X \subset R^k$).

Получающуюся задачу оптимизации можно решать с помощью общих методов численного поиска экстремума. При этом возможны два подхода. Первый из них — поиск минимума $\Psi[M]$ в пространстве элементов информационной матрицы при огра-

ничениях $M \in \mathfrak{M}$, где $\mathfrak{M}$ — множество информационных матриц: $\mathfrak{M} = \{M(\xi) | \xi \in \Xi\}$. Поскольку $\Psi[M] = \Phi[D]$ — выпуклая функция, то имеем задачу выпуклого программирования, для решения которой имеется большое число хорошо изученных численных процедур поиска оптимального решения. Второй подход — минимизация $\Psi[M(\xi)]$ по набору аргументов $\{x_i, p_i\}_{i=1}^n$ при ограничениях $x_i \in X, p_i \geq 0, \sum_{i=1}^n p_i = 1$. Эта задача не является задачей выпуклого программирования.

В обоих случаях серьезной трудностью является большая размерность экстремальной задачи. При использовании первого подхода дополнительную трудность представляют описание области $\mathfrak{M}$ и поиск плана ξ^* , соответствующего M^* (оптимальной точке из $\mathfrak{M}$), при котором необходимо решать нелинейную систему алгебраических уравнений относительно $\{p_i, x_i\}_{i=1}^n$. При втором подходе основная трудность экстремальной задачи состоит в ее многоэкстремальности.

Опишем методы, использующие специфику минимизируемых функционалов. Классический подход к задаче построения оптимальных планов

$$\xi^* = \arg \min_{\xi \in \Xi} \Phi(D(\xi))$$

состоит в следующем. Пусть имеется план ξ_s . Рассмотрим план $\xi_{s+1} = (1 - \alpha)\xi_s + \alpha\xi$ ($0 \leq \alpha \leq 1$). При достаточно малых α и необходимой гладкости функции Φ по формуле Тейлора имеем

$$\Phi[D(\xi_{s+1})] \approx \Phi[D(\xi_s)] + \alpha \Delta[D(\xi_s), D(\xi)],$$

где $\Delta(A, B)$ — производная функционала Φ в точке $D(\xi_s)$ по направлению $D(\xi) - D(\xi_s)$ (определена в п. 2.1 с заменой Ψ на Φ). Естественно выбрать план ξ таким, чтобы величина $\Delta[D(\xi_s), D(\xi)]$ была мала (но велика по модулю).

Одним из таких планов является план $\xi = \xi(x_{(s)})$ с единичной мерой, приписанной точке

$$x_{(s)} = \arg \max_{x \in X} \varphi(x, \xi_s). \quad (4.1)$$

Действительно, в классе планов вида

$$\xi_{s+1} = (1 - \alpha)\xi_s + \alpha\xi,$$

(где ξ_s — произвольный план) при малых α максимальное уменьшение функционала Φ достигается на плане

$$(1 - \alpha)\xi_s + \alpha\xi(x_{(s)}),$$

поскольку по лемме 2.1

$$\left. \frac{\partial \Phi[D(\xi_{s+1})]}{\partial \alpha} \right|_{\alpha=0+} = \text{tr } D(\xi_s) \hat{\Phi}(\xi_s) - \int \varphi(x, \xi_s) \xi_s(dx), \quad (4.2)$$

а максимальное значение $\int \varphi(x, \xi_s) \xi_s(dx)$ равно $\max_{x \in X} \varphi(x, \xi_s)$

и достигается в случае, когда мера $\xi_s(dx)$ сосредоточена в точке (4.1).

Таким образом, производная (4.2) минимальна в направлении плана $\xi(x_s)$, что соответствует локально максимальному уменьшению функционала в направлении плана $\xi(x_s)$. Если указанные действия повторять несколько раз, то получим следующий алгоритм, основной принцип которого совпадает с принципом градиентного метода — классического общего алгоритма локальной оптимизации.

Алгоритм 4.1 (процедура Федорова — Уинна).

1) Выбираем невырожденный начальный план ξ_0 вида (1.8), полагаем $s = 0$.

2) Отыскиваем точку (4.1).

3) Строим план

$$\xi_{s+1} = (1 - \gamma_s) \xi_s + \gamma_s \xi(x_s)$$

и переходим к шагу 2) с заменой s на $s + 1$.

Скорость сходимости этого алгоритма в значительной степени определяется выбором последовательности $\{\gamma_s\}$, который может быть осуществлен различными способами, аналогичными способом выбора подобных последовательностей для классического градиентного метода. Приведем два из них, наиболее известных и хорошо себя зарекомендовавших.

Первый из них —

$$\gamma_s = \arg \min_{\gamma > 0} \Phi [D((1 - \gamma) \xi_s + \gamma \xi(x_s))], \quad (4.3)$$

т. е. в качестве γ_s выбирается то значение γ , которое минимизирует значение критерия оптимальности в множестве планов вида $(1 - \gamma) \xi_s + \gamma \xi(x_s)$.

Второй способ, называемый методом деления пополам, заключается в том, что в качестве γ_s выбирается γ_{s-1} , если выполнено неравенство

$$\Phi [D((1 - \gamma_{s-1}) \xi_s + \gamma_{s-1} \xi(x_s))] < \Phi [D(\xi_s)],$$

в противном случае γ_{s-1} уменьшается вдвое до момента выполнения неравенства, после чего уменьшенное значение γ_{s-1} и выбирается в качестве γ_s .

Алгоритм 4.1 и рассмотренные ниже его модификации эффективны при малых размерностях множества $X \subset \mathbb{R}^k$ (скажем, $k < 10$). При больших размерностях X (k одного порядка с $m(m+1)/2$ или более) целесообразно искать оптимальную информационную матрицу, а потом решать систему нелинейных уравнений для определения точек и весов оптимального плана.

Несмотря на то, что сходимость алгоритма 4.1 может быть получена на основе общих теорем о сходимости методов градиентного типа, мы докажем теорему о сходимости этого алгоритма, поскольку приводимое доказательство достаточно простое и дополнительно проясняет сущность самого алгоритма.

Сначала сформулируем вспомогательное утверждение.

Лемма 4.1. Пусть выполнены условия теоремы 2.2. Тогда для любого невырожденного плана ξ

$$\max_{x \in X} \varphi(x, \xi) - \operatorname{tr} D(\xi) \bar{\Phi}(\xi) \geq \Phi[D(\xi)] - \min_{\xi} \Phi[D(\xi)]. \quad (4.4)$$

Доказательство. Рассмотрим план

$$\xi_\alpha = (1 - \alpha)\xi + \alpha\xi^*,$$

где ξ^* — Φ -оптимальный план. В силу выпуклости Φ

$$\Phi[D(\xi_\alpha)] \leq (1 - \alpha)\Phi[D(\xi)] + \alpha\Phi[D(\xi^*)],$$

$$\left. \frac{\partial \Phi[D(\xi_\alpha)]}{\partial \alpha} \right|_{\alpha=0+} \leq \Phi[D(\xi^*)] - \Phi[D(\xi)].$$

С другой стороны, из формулы (4.2) для $\xi_s = \xi$, $\xi_* = \xi^*$ имеем

$$\begin{aligned} \left. \frac{\partial \Phi[D(\xi_\alpha)]}{\partial \alpha} \right|_{\alpha=0+} &= \operatorname{tr} D(\xi) \bar{\Phi}(\xi) - \int_X \varphi(x, \xi) \xi^*(dx) \geq \\ &\geq \operatorname{tr} D(\xi) \bar{\Phi}(\xi) - \max_{x \in X} \varphi(x, \xi). \end{aligned}$$

Объединяя эти неравенства, получаем (4.4). Лемма доказана.

Из утверждения леммы 4.1 следует, в частности, простое правило прекращения счета в алгоритме 4.1 (то же самое правило можно использовать и при использовании приведенных ниже алгоритмов 4.2, 4.3).

Пусть задана необходимая точность конечного результата

$$\Phi[D(\xi_s)] - \inf_{\xi \in \mathfrak{B}} \Phi[D(\xi)] \leq \nu, \quad \nu > 0.$$

В алгоритме 4.1 расчеты следует прекратить, как только на шаге 2) будет выполнено неравенство

$$\max_{x \in X} \varphi(x, \xi_s) - \operatorname{tr} D(\xi_s) \bar{\Phi}(\xi_s) \leq \nu. \quad (4.5)$$

Из (4.4) следует, что если выполнено второе из этих двух неравенств, то выполнено и первое.

Положим $\mathcal{D}(C) = \{D | \Phi[D] \leq C, D^{-1} \in \mathfrak{M}\}$. Тогда имеет место следующее утверждение о сходимости алгоритма 4.1.

Теорема 4.1. Пусть множества X и $\mathfrak{B}$ — компакты, Φ — такой выпуклый дифференцируемый функционал, что Φ -оптимальный план ξ^* невырожден и существуют ограниченные производные

$$\partial^2 \Phi[D] / (\partial D_{ij} \partial D_{uv}), \quad i, j, u, v = 1, 2, \dots, m, \quad (4.6)$$

по элементам матрицы $D \in \mathcal{D}(C)$ при любом $C > 0$. Пусть, далее, последовательность $\{\gamma_s\}$ определяется по формуле (4.3).

Тогда алгоритм 4.1 сходится в том смысле, что

$$\lim_{s \rightarrow \infty} \Phi[D(\xi_s)] = \min_{\xi \in \mathfrak{B}} \Phi[D(\xi)],$$

причем из последовательности $\{\xi_s\}$ можно выделить подпоследовательность планов, слабо сходящуюся к одному из Φ -оптимальных планов ξ^* . Если, кроме того, функционал Φ строго выпуклый, то

$$\lim_{s \rightarrow \infty} D(\xi_s) = D(\xi^*).$$

Доказательство. В силу (4.3) алгоритм 4.1 релаксационный, т. е. последовательность $\{\Phi[D(\xi_s)]\}$ монотонно убывающая. Поскольку эта последовательность ограничена снизу, то она сходится, т. е.

$$\lim_{s \rightarrow \infty} \Phi[D(\xi_s)] = \Phi[D(\xi^*)],$$

где ξ^* — некоторый план, существующий в силу компактности множества $\mathfrak{M}$.

Покажем, что план ξ^* оптимален. Предположим противное:

$$\Phi[D(\xi^*)] > \min_{\xi \in \mathfrak{E}} \Phi[D(\xi)].$$

Тогда в силу леммы 4.1 и сходимости последовательности $\{\Phi[D(\xi_s)]\}$ можно найти такие $\delta > 0$ и s_0 , что для всех $s \geq s_0$ выполняется неравенство

$$\max_{x \in X} \varphi(x, \xi_s) - \text{tr } D(\xi_s) \Phi(\xi_s) \geq \delta > 0. \quad (4.7)$$

В силу дважды дифференцируемости Φ и формулы Тейлора величины $\Phi[D(\xi_{s+1})]$ могут быть представлены в виде

$$\Phi[D(\xi_{s+1})] = \Phi[D(\xi_s)] - \gamma_s [\varphi(x_{s+1}, \xi_s) - \text{tr } D(\xi_s) \Phi(\xi_s)] - \gamma_s^2 f''(\bar{\gamma}), \quad (4.8)$$

где

$$f(\gamma) = \Phi[D[(1-\gamma)\xi_s + \gamma\xi(x_s)] - \Phi[D(\xi_s)], \\ 0 \leq \bar{\gamma} \leq \gamma_s.$$

Положим $C = \Phi[D(\xi_0)]$. Из ограниченности производных (4.6) на множестве $\mathcal{D}(C)$ следует, что существует такое $Q < \infty$ (зависящее от C), при котором имеет место неравенство

$$0 \leq \sup_{\gamma \in [0, 1]} f''(\gamma) \leq Q. \quad (4.9)$$

Из (4.7) — (4.9) получаем

$$\Phi[D(\xi_s)] - \Phi[D(\xi_{s+1})] \geq \gamma_s \delta. \quad (4.10)$$

Покажем теперь, что при $s \geq s_0$ должно выполняться неравенство $\gamma_s \geq \delta/Q$. Имеем

$$f'(0) = \varphi(x_s, \xi_s) - \text{tr } D(\xi_s) \Phi(\xi_s) \geq \delta.$$

Поскольку для любых $\gamma_1, \gamma_2 \in [0, 1]$

$$|f'(\gamma_1) - f'(\gamma_2)| \leq |\gamma_1 - \gamma_2| \max_{0 \leq \gamma \leq 1} f''(\gamma),$$

то при $\gamma_1 = 0, \gamma_2 = \gamma$ получим

$$|f'(0) - f'(\gamma)| \leq \gamma Q;$$

откуда следует, что при

$$\gamma < |f'(0)|/Q \leq \delta/Q$$

производная $f'(\gamma)$ в нуль не обращается, и поэтому при $\gamma < \delta/Q$ функция $f(\gamma)$ минимума не достигает.

Из (4.10) имеем

$$\Phi[D(\xi_s)] - \Phi[D(\xi_{s+1})] \geq \delta^2/Q.$$

Отсюда получаем

$$\lim_{s \rightarrow \infty} \Phi[D(\xi_s)] \leq C - \sum_{s=0}^{\infty} \frac{\delta^2}{Q} = -\infty.$$

Это противоречит тому, что

$$\inf \Phi[D(\xi)] > -\infty.$$

Следовательно, план ξ^* оптимальный.

Последнее утверждение теоремы следует из слабой компактности множества Ξ (см. с. 91) и из теоремы 1.4. Теорема доказана.

Приведем две модификации алгоритма 4.1. Численные расчеты показывают, что скорость сходимости алгоритма 4.1 увеличивается, если в качестве допустимых включить движения по направлениям, определяемым точками x_{is} плана ξ_s с отрицательными γ_s .

Алгоритм 4.2.

1) Выбираем невырожденный начальный план ξ_0 , полагаем $s = 0$.

2) Отыскиваем точку

$$x_s = \arg \max \{ \varphi(x_s^+, \xi_s) - \text{tr } D(\xi_s) \Phi(\xi_s), \text{tr } D(\xi_s) \Phi(\xi_s) - \varphi(x_s^-, \xi_s) \},$$

где

$$x_s^+ = \arg \max_{x \in X} \varphi(x, \xi_s), \quad x_s^- = \arg \min_{x \in X_s} \varphi(x, \xi_s),$$

X_s — множество точек плана ξ_s .

3) Строим план $\xi_{s+1} = (1 - \beta_s)\xi_s + \beta_s \xi(x_s)$, где $\beta_s = \gamma_s$, если $x_s = x_s^+$, и $\beta_s = -\min\{\alpha_s, \rho_{is}/(1 - \rho_{is})\}$, если $x_s = x_s^-$ (ρ_{is} — вес точки $x_{is} = x_s^-$ в плане ξ_s); последовательность $\{\alpha_s\}$ выбирается по тем же правилам, что и $\{\gamma_s\}$.

4) Переходим к шагу 2) с заменой s на $s + 1$.

При использовании алгоритма 4.2 удается исключить те точки начального плана, которые выбраны неудачно. Кроме

того, алгоритм 4.2 выгоднее алгоритма 4.1 тем, что при его использовании в плане, приближенном к оптимальному, содержится сравнительно мало точек (всем «плохим» точкам приписываются нулевые веса и они тем самым исключаются из плана). С вычислительной точки зрения алгоритм 4.2 имеет тот недостаток, что на каждой его итерации решается довольно сложная задача — ищется максимум функции $\varphi(x, \xi_s)$ на X , а результат решения этой задачи — точка x_s^+ — используется не всегда. В случае большой размерности множества X более эффективен следующий алгоритм, в котором экстремальная задача (1.1) решается только в том случае, когда не удастся добиться гарантированного уменьшения функционала значительно более простыми операциями уменьшения и добавления весов точкам плана.

Алгоритм 4.3.

1) Имеем невырожденный план

$$\xi_s = \left\{ \begin{matrix} x_1, & \dots, & x_n \\ p_1, & \dots, & p_n \end{matrix} \right\}.$$

Отыскиваем

$$x_{I_s} = x_s^- = \arg \min_{x \in X_s} \varphi(x, \xi_s).$$

2) Если

$$\varphi(x_s^-, \xi_s) < \text{tr } D(\xi_s) \bar{\Phi}(\xi_s) - \delta_s,$$

то составляем план

$$\xi_{s+1} = (1 + \beta_s) \xi_s - \beta_s \xi(x_s^-),$$

где последовательность $\{\beta_s\}$ выбирается по тем же правилам, что и $\{\gamma_s\}$ в алгоритме 4.1, но при условии $\beta_s \leq p_{I_s}/(1 - p_{I_s})$, и переходим к шагу 6).

3) Отыскиваем

$$x_{I_s} = \arg \max_{x \in X_s} \varphi(x, \xi_s);$$

если

$$\varphi(x_{I_s}) > \text{tr } D(\xi_s) \bar{\Phi}(\xi_s) + \delta_s,$$

то переходим к шагу 5) с $x_{(s)} = x_{I_s}$, в противном случае — к шагу 4).

4) Отыскиваем точку (4.1).

5) Составляем план

$$\xi_{s+1} = (1 - \gamma_s) \xi_s + \gamma_s \xi(x_{(s)}).$$

6) Переходим к шагу 1) с заменой s на $s + 1$.

Последовательность $\{\delta_s\}$ определяет минимальную точность, получаемую на каждом шаге алгоритма. Для обеспечения сходимости она должна выбираться с учетом условия

$$\delta_s \geq c \left[\max_{x \in X} \varphi(x, \xi_s) - \text{tr } D(\xi_s) \bar{\Phi}(\xi_s) \right], \quad c > 0.$$

В частности, хорошо зарекомендовал себя следующий выбор $\{\delta_s\}$:

$$\delta_s = \min \left\{ \frac{1}{2} (\mu_{s-1} - \operatorname{tr} D(\xi_{s-1}) \bar{\Phi}(\xi_{s-1})), \delta_{s-1} \right\},$$

если на $(s-1)$ -м шаге вычислялся $\mu_{s-1} = \max_{x \in X} \varphi(x, \xi_{s-1})$, и $\delta_s = \delta_{s-1}$, если не вычислялся.

Доказательства сходимости алгоритмов 4.2 и 4.3 аналогичны доказательству сходимости алгоритма 4.1.

Рассмотрим особенности приведенных алгоритмов для D -критерия. При D -оптимальном планировании

$$\varphi(x, \xi) = d(x, \xi) = f^T(x) D(\xi) f(x), \quad \operatorname{tr} D(\xi) \bar{\Phi}(\xi) = m$$

для любых $x \in X$ и $\xi \in \Xi$. Специфика D -критерия позволяет существенно упростить и вспомогательные вычисления, которые необходимо проводить при использовании алгоритмов 4.1—4.3.

Теорема 4.2. Положим

$$\xi_{s+1} = (1 - \gamma) \xi_s + \gamma \xi(x),$$

где план ξ_s невырожден, план $\xi(x)$ сосредоточен в точке $x \in X$, $-p(x)/[1 - p(x)] < \gamma < 1$, $p(x)$ — вес точки x в плане ξ_s ($p(x) = 0$, если точка x в план ξ_s не входит). Тогда

$$a) \quad D(\xi_{s+1}) = \frac{1}{1-\gamma} \left[I_m - \frac{\gamma D(\xi_s) f(x) f^T(x)}{1-\gamma + \gamma d(x, \xi_s)} \right] D(\xi_s); \quad (4.11)$$

$$б) \quad \det M(\xi_{s+1}) = (1 - \gamma)^m \left[1 + \frac{\gamma d(x, \xi_s)}{1-\gamma} \right] \det M(\xi_s); \quad (4.12)$$

в) если $\gamma = \gamma_s$ выбирается из условия (4.3), $\Phi[D] = \ln \det D$ и $x = x_s$, то

$$\gamma = \gamma_s = \frac{d(x, \xi_s) - m}{(d(x, \xi_s) - 1)m}; \quad (4.13)$$

г) максимальное увеличение определителя (4.12) в случае $\gamma > 0$ достигается при $\gamma = \gamma_s$, выбранном по (4.13), при этом

$$\det M((1 - \gamma_s) \xi_s + \gamma_s \xi(x)) = \left(\frac{m-1}{d(x, \xi_s) - 1} \right)^{m-1} \left(\frac{d(x, \xi_s)}{m} \right)^m \det M(\xi_s). \quad (4.14)$$

Доказательство. а) Имеем

$$M(\xi_{s+1}) = (1 - \gamma) M(\xi_s) + \gamma M(\xi(x)) = (1 - \gamma) M(\xi_s) + \gamma f(x) f^T(x),$$

откуда

$$D(\xi_{s+1}) = M^{-1}(\xi_{s+1}) = [(1 - \gamma) M(\xi_s) + \gamma f(x) f^T(x)]^{-1} = \\ = (1 - \gamma)^{-1} \left[I_m + \frac{\gamma}{1-\gamma} D(\xi_s) f(x) f^T(x) \right]^{-1} D(\xi_s).$$

Положим

$$\frac{\gamma}{1-\gamma} D(\xi_s) f(x) = A, \quad f^T(x) = B, \quad p = m, \quad q = 1$$

и воспользуемся формулой (1.10) из приложения 1. Получим

$$\begin{aligned} \left[I_m + \frac{\gamma}{1-\gamma} D(\xi_s) f(x) f^T(x) \right]^{-1} &= [I_p + AB]^{-1} = \\ &= I_p - A(I_q + BA)^{-1} B = I_m - \frac{\gamma}{1-\gamma} D(\xi_s) f(x) \times \\ &\times \left[1 + \frac{\gamma}{1-\gamma} f^T(x) D(\xi_s) f(x) \right]^{-1} f^T(x). \end{aligned}$$

откуда и следует (4.11).

б) Для того чтобы получить (4.12), запишем

$$\begin{aligned} M(\xi_{s+1}) &= (1-\gamma) \left[M(\xi_s) + \frac{\gamma}{1-\gamma} f(x) f^T(x) \right], \\ \det M(\xi_{s+1}) &= (1-\gamma)^n \det \left[M(\xi_s) + \frac{\gamma}{1-\gamma} f(x) f^T(x) \right] \end{aligned}$$

и воспользуемся формулой (1.12) из приложения 1 с

$$A = M(\xi_{s+1}), \quad B = \sqrt{\frac{\gamma}{(1-\gamma)}} f(x).$$

в) Для того чтобы получить (4.13), прологарифмируем (4.12), продифференцируем полученное выражение по γ и производную приравняем нулю. Получим

$$\frac{d(x, \xi_s) - 1}{1 - \gamma + \gamma d(x, \xi_s)} - \frac{m-1}{1-\gamma} = 0,$$

что равносильно утверждению в).

Утверждение г) — простое следствие б) и в). Теорема доказана.

Формула (4.11) позволяет вычислять дисперсионную матрицу плана ξ_{s+1} при всех s , не прибегая к операции обращения матриц; операцию обращения, таким образом, необходимо применять только для вычисления $D(\xi_0)$. Польза этого несомненна не только в случае D -критерия, но и для других критериев Ф.

4.2. Методы построения дискретных оптимальных планов. Непрерывные планы с точки зрения практического эксперимента являются приближенным решением исходной задачи планирования. Приближение тем лучше, чем больше число N возможных наблюдений (меньше сказывается дискретность весов $p_i = r_i/N$). Естественной процедурой построения планов, приближенных к дискретному оптимальному

$$\xi_N^* = \arg \max_{\xi_N \in \Xi} \Psi[M(\xi_N)], \quad (4.15)$$

является процедура округления непрерывного Ψ -оптимального плана

$$\xi^* = \left\{ \begin{matrix} x_1^*, \dots, x_n^* \\ p_1^*, \dots, p_n^* \end{matrix} \right\}, \quad (4.16)$$

Простейшие процедуры округления плана (4.16) состоят в том, что в качестве плана, приближенного к (4.15), выбирается

$$\tilde{x}_N = \left\{ \frac{x_1^* + \alpha_1}{N}, \dots, \frac{x_n^* + \alpha_n}{N} \right\}, \quad (4.17)$$

где целые числа r_i ($i = 1, \dots, n$) должны удовлетворять соотношениям $r_i \leq N p_i^*$ и могут быть определены, например, по формуле

$$r_i = r_i(N) = \lceil (N - n) p_i^* \rceil \quad (4.18)$$

или

$$r_i = r_i(N) = \lfloor N p_i^* \rfloor \quad (4.19)$$

(здесь $\lfloor a \rfloor$ — целая часть a , $\lceil a \rceil$ — ближайшее к a целое число, которое больше a ; целые неотрицательные числа α_i ($i = 1, \dots, n$) выбираются произвольно с учетом ограничения

$$\sum_{i=1}^n \alpha_i = N_1 = N - \sum_{i=1}^n r_i, \quad N_1 \leq n.$$

Очевидной представляется следующая процедура выбора $\alpha_1, \dots, \alpha_n$: величины $N p_i^* - r_i$ располагаются в порядке убывания их значений; если $N p_i^* - r_i$ стоит на одном из первых N_1 мест в указанном упорядоченном наборе, то полагаем $\alpha_i = 1$, в противном случае $\alpha_i = 0$.

Оценки точности обеих процедур округления оказываются одинаковыми в том смысле, что точность зависит только от близости числа N_1 к N . Само же число N_1 при использовании разных процедур имеет различные значения.

Теорема 4.3. Пусть выполнены условия монотонности и однородности для функционала Ψ . Тогда для плана (4.17) при любом способе получения чисел r_i ($r_i \leq N p_i^*$, $i = 1, \dots, n$) справедливо неравенство

$$\Psi[M(\xi^*)] - \Psi[M(\tilde{x}_N)] \leq \left(1 - \frac{\gamma(N - N_1)}{\gamma(N)}\right) \Psi[M(\xi^*)], \quad (4.20)$$

где $N_1 = N - \sum_{i=1}^n r_i$.

Доказательство. Имеем

$$\begin{aligned} \gamma(N) \Psi[M(\tilde{x}_N)] &= \Psi[NM(\tilde{x}_N)] = \\ &= \Psi \left[\sum_{i=1}^n r_i f(x_i^*) f^T(x_i^*) + \sum_{i=1}^n \alpha_i f(x_i^*) f^T(x_i^*) \right] \geq \\ &\geq \Psi \left[\sum_{i=1}^n N p_i^* f(x_i^*) f^T(x_i^*) \right] = \Psi[(N - N_1) M(\xi^*)] = \\ &= \gamma(N - N_1) \Psi[M(\xi^*)], \end{aligned}$$

что равносильно (4.20). Теорема доказана.

Из (4.20) с учетом того, что $N_1 \leq n$ и $\Psi[M(\xi_N^*)] \leq \Psi[M(\xi^*)]$, вытекают следующие неравенства:

$$\Psi[M(\xi^*)] - \Psi[M(\bar{\xi}_N)] \leq \left(1 - \frac{\gamma(N-n)}{\gamma(N)}\right) \Psi[M(\xi^*)], \quad (4.21)$$

$$\Psi[M(\xi^*)] - \Psi[M(\xi_N^*)] \leq \left(1 - \frac{\gamma(N-n)}{\gamma(N)}\right) \Psi[M(\xi^*)] \quad (4.22)$$

(оценка близости значений критерия для непрерывного Ψ -оптимального и дискретного Ψ -оптимального планов),

$$\Psi[M(\xi_N^*)] - \Psi[M(\bar{\xi}_N)] \leq \left(1 - \frac{\gamma(N-n)}{\gamma(N)}\right) \Psi[M(\xi^*)] \quad (4.23)$$

(оценка близости значений критерия планов (4.17) и (4.15)).

Из рассмотрений оценок (4.20)–(4.23) можно сделать следующие выводы:

а) вообще говоря, выгоднее округлять те непрерывные планы (4.16), у которых количество точек n в плане мало;

б) при фиксированном N из двух процедур округления (4.18), (4.19) выгоднее та, для которой величина $N_1 = N - \sum_{i=1}^n r_i$ меньше (процедура гарантирует большую близость приближенного плана к оптимальному);

в) более выгодными, чем процедуры (4.18), (4.19), могут оказаться их модификации, основанные на замене в формулах (4.18), (4.19) N на

$$\bar{N} = \min \left\{ k \geq N \mid \sum_{i=1}^n r_i(k) = N \right\}$$

и составлении плана (4.17) с $\alpha_i = 0$ ($i = 1, \dots, n$). Оценка (4.20) в этом случае имеет вид

$$\Psi[M(\xi^*)] - \Psi[M(\bar{\xi}_N)] \leq \left(1 - \frac{\gamma(N)}{\gamma(\bar{N})}\right) \Psi[M(\xi^*)];$$

г) если N мало (не намного больше n), то приведенные процедуры неэффективны, и более целесообразно для приближенного построения дискретных оптимальных планов использовать другие алгоритмы (в частности, рассматриваемый ниже).

Переходим к рассмотрению алгоритма, идея которого совпадает с идеей алгоритма 4.2 и в основе которого (как и алгоритмов 4.1–4.3) лежит теорема эквивалентности. Для удобства изложение снова проводится в терминах функционалов Φ от дисперсионных матриц и в предположении, что дискретные оптимальные планы

$$\xi_N^* = \arg \min_{\xi_N \in E_N} \Phi[D(\xi_N)] \quad (4.24)$$

невырождены.

Алгоритм 4.4.

1) Выбираем невырожденный начальный план

$$\xi_N^{(0)} = \left\{ x_1^{(0)}, \dots, x_N^{(0)} \right\}_{1/N, \dots, 1/N}$$

и полагаем $s = 0$.

2) Отыскиваем точку

$$x_{(s)}^+ = \arg \max_{x \in X} \Phi(x, \xi_N^{(s)}). \quad (4.25)$$

3) Составляем план

$$\begin{aligned} \xi_{N+1}^{(s)} &= \left(1 - \frac{1}{N+1}\right) \xi_N^{(s)} + \frac{1}{N+1} \xi(x_{(s)}^+) = \\ &= \left\{ \begin{array}{cccc} x_1^{(s)} & , & \dots & , & x_N^{(s)} & , & x_{(s)}^+ \\ 1/(N+1), & \dots, & 1/(N+1), & 1/(N+1) \end{array} \right\}. \end{aligned}$$

4) Отыскиваем точку

$$x_{(s)}^- = \arg \min_{x \in X_s} \Phi(x, \xi_{N+1}^{(s)}),$$

где X_s — множество точек плана $\xi_{N+1}^{(s)}$.

5) Составляем план

$$\xi_N^{(s+1)} = \left(1 + \frac{1}{N}\right) \xi_{N+1}^{(s)} - \frac{1}{N} \xi(x_{(s)}^-) = \left\{ \begin{array}{ccc} x_1^{(s+1)}, & \dots, & x_N^{(s+1)} \\ 1/N, & \dots, & 1/N \end{array} \right\}.$$

6) Если $x_{(s)}^-$ совпала с $x_{(s)}^+$, то вычисления прекращаем (дальнейшего уменьшения функционала Φ происходить не будет), в противном случае — переходим к шагу 2) с заменой s на $s+1$.

Суть алгоритма 4.4 чрезвычайно проста — на каждой итерации в план добавляется «лучшая» точка из X , а затем из полученного плана удаляется «худшая» его точка. Исходя из этих соображений, сходимость последовательности $\{\Phi[D(\xi_N^{(s)})]\}$ очевидна, однако ее предел не обязательно совпадает с

$$\inf_{\xi_N \in \Xi_N} \Phi[D(\xi_N)].$$

Гарантировать это совпадение можно лишь при совпадении плана с одним из непрерывных оптимальных планов (что определяется при выполнении шага 2) алгоритма 4.4).

Легко построить различные модификации алгоритма 4.4. Например, на каждой итерации можно включать в план и исключать из него не по одной, а по несколько точек.

4.3. О построении планов, оптимальных в $\Xi(n)$. Для построения планов, приближенных к оптимальным в множестве $\Xi(n)$ (определение этого множества содержится в п. 1.6), можно использовать очевидную модификацию алгоритма 4.4. Планы $\xi^{(s)}$

должны быть вида $\xi^{(s)} = \xi^{(s)}(n)$, т. е. вида (1.8) с фиксированным n ; на шаге 3) вес плана $\xi(x_{(s)}^+)$ может быть не равным $1/(N+1)$, а выбираться любым из способов, рассматривавшихся при описании алгоритма 4.1; на шаге 5) нужно составлять план

$$\xi^{(s+1)} = \frac{1}{1-\rho_-} \xi^{(s)} - \frac{\rho_-}{1-\rho_-} \xi(x_{(s)}^-),$$

где ρ_- — вес точки $x_{(s)}^-$ в плане $\xi^{(s)}$ (для того чтобы удалить точку $x_{(s)}^-$ из плана $\xi^{(s)}$).

Упражнения.

1. Предположим, что для функции регрессии $\eta(x) = \theta_1 + \theta_2 x$ на отрезке $[-1, 1]$ требуется численно построить D -оптимальный план. Выбирая γ , по формуле (4.13) и начальный план

$$\xi_0 = \begin{Bmatrix} -1/2 & 1/2 \\ 1/2 & 1/2 \end{Bmatrix},$$

провести по пять итераций алгоритмов 4.1, 4.2, 4.3.

2. Округлить различными способами D -оптимальный план (3.14) для квадратичной регрессии на отрезке $[-1, 1]$ при $N = 10$ и $N = 11$. На этих примерах проверить точность формулы (4.20).

Глава 4

ПЛАНИРОВАНИЕ РЕГРЕССИОННОГО ЭКСПЕРИМЕНТА (ДАЛЬНЕЙШИЕ ПОСТАНОВКИ ЗАДАЧ И РЕЗУЛЬТАТЫ)

В данной главе содержатся методы оптимального планирования регрессионного эксперимента: с одной стороны, для ряда частных постановок, которые могут быть вписаны в рамки классической, и, с другой — для схем регрессии, отличных от рассмотренной в гл. 3. Последние дают некоторое представление о современных направлениях развития математической теории планирования регрессионного эксперимента. При первом чтении § 2—4 могут быть пропущены.

§ 1. Планы первого порядка

1.1. Основные понятия. Основные результаты данного параграфа можно получить из общих теорем гл. 3. Представляется поучительным, однако, получение этих важных для приложений результатов непосредственно из определений и элементарных фактов алгебры и анализа.

Предположим, что функция регрессии — полином первой степени от m переменных $x_1, \dots, x_m$, т. е.

$$\eta(x) = \theta_0 + \theta_1 x_1 + \dots + \theta_m x_m, \quad (1.1)$$

а ошибки измерений, как обычно, центрированы, некоррелированы и имеют одинаковую дисперсию σ^2 .

Для удобства в данном параграфе нумерация неизвестных параметров начинается с нуля (число неизвестных параметров равно $m + 1$). Поэтому информационные и дисперсионные матрицы планов имеют размер $(m + 1) \times (m + 1)$.

Дискретный план

$$\xi_N = \left\{ \begin{matrix} x_1, & \dots, & x_N \\ 1/N, & \dots, & 1/N \end{matrix} \right\}, \quad (1.2)$$

который позволяет получить несмещенные МНК-оценки параметров $\theta_0, \dots, \theta_m$ функции регрессии (1.1), будем называть *планом первого порядка*.

Планы первого порядка находят широкое применение в практике экспериментальных исследований. Это связано, в частности, с тем, что любая линейная по неизвестным параметрам функция регрессии может быть приведена к виду (1.1) введением новых переменных, и с тем, что построение планов первого порядка является важной частью некоторых стратегий планирования экстремального эксперимента (см. гл. 6).

План (1.2) удобно представлять в виде матрицы

$$D = \begin{bmatrix} x_{11} & x_{12} & \dots & x_{1m} \\ x_{21} & x_{22} & \dots & x_{2m} \\ \dots & \dots & \dots & \dots \\ x_{N1} & x_{N2} & \dots & x_{Nm} \end{bmatrix}, \quad (1.3)$$

где x_{ji} ($i = 1, \dots, m$) — координаты точки x_j ($j = 1, \dots, N$). Матрица плана (в терминологии гл. 1) имеет вид

$$F = \begin{bmatrix} x_{10} & x_{11} & \dots & x_{1m} \\ x_{20} & x_{21} & \dots & x_{2m} \\ \dots & \dots & \dots & \dots \\ x_{N0} & x_{N1} & \dots & x_{Nm} \end{bmatrix}, \quad (1.4)$$

где $x_{j0} = 1$ ($j = 1, 2, \dots, N$).

Поскольку план (1.2) предполагается невырожденным, то

$$N \geq m + 1, \quad \text{rang } D = m, \quad \text{rang } F = m + 1.$$

В соответствии с определением из § 1 гл. 2 план первого порядка (1.2) ортогональный, если столбцы матрицы F попарно ортогональны, т. е. если

$$\sum_{i=1}^N x_{it} x_{ij} = 0, \quad i, j = 0, 1, \dots, m, \quad i \neq j. \quad (1.5)$$

Для ортогонального плана первого порядка выполняется условие симметрии

$$\sum_{i=1}^N x_{it} = 0, \quad i = 1, 2, \dots, m, \quad (1.6)$$

которое следует из (1.5) при $j = 0$.

На двух примерах посмотрим, как влияет выбор плана на дисперсии $D\hat{\theta}_j$, МНК-оценок $\hat{\theta}_j$ неизвестных параметров θ_j ($j = 0, 1, \dots, m$).

Пример 1.1. Пусть $N = 6$, $m = 2$. Рассмотрим два плана $\xi_6^{(1)}$ и $\xi_6^{(2)}$ вида (1.2), которым соответствуют матрицы

$$F_1 = \begin{bmatrix} 1 & -1 & 0 \\ 1 & 1 & -\sqrt{3} \\ 1 & -1 & 0 \\ 1 & 1 & -\sqrt{3} \\ 1 & -1 & 0 \\ 1 & 1 & 0 \end{bmatrix} = (x_{it}^{(1)}), \quad F_2 = \begin{bmatrix} 1 & -1 & -1 \\ 1 & -1 & -1 \\ 1 & 1 & -1 \\ 1 & 1 & -1 \\ 1 & -1 & 1 \\ 1 & -1 & 1 \end{bmatrix} = (x_{it}^{(2)}).$$

Очевидно, что план $\xi_6^{(1)}$ — ортогональный план первого порядка и

$$\sum_{i=1}^6 (x_{ij}^{(1)})^2 = 6, \quad j = 0, 1, 2.$$

Поэтому дисперсии МНК-оценок параметров $\theta_0, \theta_1, \theta_2$ равны $\sigma^2/6$, и (в силу ортогональности $\xi_6^{(1)}$) эти оценки некоррелированы.

Для плана $\xi_6^{(2)}$ также

$$\sum_{i=1}^6 (x_{ij}^{(2)})^2 = 6, \quad j = 0, 1, 2.$$

План $\xi_6^{(2)}$ не является ортогональным, поскольку дисперсионная матрица МНК-оценок равна

$$\sigma^2 (F_2^T F_2)^{-1} = \sigma^2 \begin{bmatrix} 1/4 & 1/8 & 1/8 \\ 1/8 & 1/4 & 1/8 \\ 1/8 & 1/8 & 1/4 \end{bmatrix}.$$

Отсюда получаем, что дисперсии всех трех МНК-оценок равны $\sigma^2/4$, т. е. больше, чем для плана $\xi_6^{(1)}$. Кроме этого, МНК-оценки при использовании плана $\xi_6^{(2)}$ получаются коррелированными.

Пример 1.2. Пусть $m = 3, N = 4$,

$$F_1 = \begin{bmatrix} 1 & -1 & -1 & 1 \\ 1 & 1 & -1 & -1 \\ 1 & -1 & 1 & -1 \\ 1 & 1 & 1 & 1 \end{bmatrix}, \quad F_2 = \begin{bmatrix} 1 & -1 & -1 & -1 \\ 1 & 1 & -1 & -1 \\ 1 & -1 & 1 & -1 \\ 1 & 1 & 1 & 1 \end{bmatrix}.$$

План $\xi_4^{(1)}$, соответствующий матрице F_1 , является ортогональным планом первого порядка, при использовании которого получаем $D\hat{\theta}_j = \sigma^2/4$ ($j = 0, 1, 2, 3$). Легко проверить, что при использовании неортогонального плана $\xi_4^{(2)}$, соответствующего матрице F_2 ,

$$D\hat{\theta}_0 = \sigma^2, \quad D\hat{\theta}_i = \sigma^2/2, \quad i = 1, 2, 3,$$

и, следовательно, план $\xi_4^{(2)}$ существенно хуже плана $\xi_4^{(1)}$ с точки зрения величин дисперсий МНК-оценок неизвестных параметров.

В приведенных примерах при использовании ортогональных планов дисперсии МНК-оценок имели меньшие значения, чем неортогональные. В п. 1.2 показано, что это было не случайно: при определенных предположениях дискретные А-оптимальные (т. е. такие, для которых $\sum_{j=0}^m D\hat{\theta}_j$ минимальна) планы первого порядка сосредоточены в множестве ортогональных планов.

1.2. А-оптимальность ортогональных планов первого порядка. Как и в п. 1.1 рассматривается линейная регрессионная модель

$$(Y, F\theta, \sigma^2/N), \quad (1.7)$$

где матрица F имеет вид (1.4). Поскольку при проведении измерений значения факторов (т. е. переменных $x_1, \dots, x_m$) ограничены (множество планирования X здесь не конкретизируется), будем считать, что выполнены ограничения

$$\sum_{i=1}^N x_{ii}^2 = C_i^2, \quad i = 1, \dots, m, \quad (1.8)$$

где C_i^2 заданы. Заметим, что если выполнено (1.8), то выполнено и

$$\sum_{i=1}^N x_{ii}^2 = C_i^2, \quad i = 0, 1, \dots, m, \quad (1.9)$$

где $C_0^2 = N$.

Теорема 1.1. Пусть заданы модель (1.7) и класс дискретных планов первого порядка вида (1.2) таких, что выполнено

$$\sum_{i=1}^N x_{ii}^2 \leq C_i^2, \quad i = 0, 1, \dots, m. \quad (1.10)$$

Тогда

$$D\hat{\theta}_i \geq \sigma^2/C_i^2, \quad j = 0, 1, \dots, m. \quad (1.11)$$

Причем в том случае, когда план ортогональный и выполнено (1.9), в (1.11) достигается знак равенства, и поэтому любой ортогональный план первого порядка, удовлетворяющий (1.9), является A -оптимальным в рассматриваемом множестве планов.

Доказательство. В силу невырожденности планов матрица F имеет ранг $m+1$, поэтому матрица $M = F^T F$ невырожденная и, следовательно, положительно определенная. В силу положительной определенности матрицы M ее можно представить в виде $M = B^T B$, где B — невырожденная верхняя треугольная матрица (см. теорему 1.4 из приложения 1).

Пусть

$$F^T F = B^T B, \quad (F^T F)^{-1} = (B^T B)^{-1} = B^{-1} (B^{-1})^T. \quad (1.12)$$

Примем обозначения

$$F^T F = (m_{ij}), \quad B = (b_{ij}), \quad B^{-1} = (\bar{b}_{ij}), \quad (F^T F)^{-1} = (d_{ij}).$$

Очевидно, что B^{-1} — тоже верхняя треугольная матрица, причем $\bar{b}_{ii} = 1/b_{ii}$ ($i = 0, 1, \dots, m$). Из формулы для дисперсий МНК-оценок $\hat{\theta}_j$ следует, что

$$D\hat{\theta}_j = \sigma^2 d_{jj}, \quad j = 0, 1, \dots, m.$$

Из (1.12) вытекает равенство

$$d_{jj} = \sum_{i=0}^m \bar{b}_{ii}^2, \quad j = 0, 1, \dots, m.$$

Имеют место очевидные неравенства:

$$\sum_{i=0}^m \bar{b}_{ii}^2 \geq \bar{b}_{ii}^2 = 1/b_{ii}^2 \geq 1/\sum_{i=0}^m b_{ii}^2.$$

Но в силу (1.12) и (1.10)

$$\sum_{i=0}^m b_{ii}^2 = m_{ii} = \sum_{i=1}^N x_{ii}^2 \leq C_j^2,$$

и поэтому $d_{ii} \geq C_j^{-2}$, откуда и следует (1.11).

Если план (1.2) ортогональный и выполнено (1.9), то все неравенства превращаются в равенства. Теорема доказана.

Очевидно, что если в (1.8) все C_j ($j=1, \dots, m$) равны, то соответствующий ортогональный план является и ротатабельным. Такие планы удобны при практическом использовании, поскольку: очень просто вычисляются оценки неизвестных параметров; оценки параметров некоррелированы (т. е. параметры регрессии оцениваются независимо друг от друга); все параметры регрессии определяются с одинаковой и минимальной (в смысле теоремы 1.1) дисперсией; планирование является ротатабельным (т. е. информация, содержащаяся в уравнении регрессии, равномерно «размазана» по сфере).

Особенно часто ортогональные ротатабельные планы первого порядка используются при планировании экстремальных экспериментов (см. гл. 6), когда оптимизация производится без использования ЭВМ.

В п. 1.3 приведен пример использования теоремы 1.1 в задачах другого рода.

1.3. Задача оптимального взвешивания. Пусть $\beta_1, \beta_2, \beta_3$ — неизвестные веса соответственно трех предметов A_1, A_2, A_3 . Необходимо наилучшим образом оценить эти веса по результатам четырех взвешиваний на одночашечных весах, которые взвешивают предметы со случайной ошибкой.

Введем величины z_{j1}, z_{j2}, z_{j3} . Будем считать, что если предмет A_i находится на весах при j -м измерении, то $z_{ji} = 1$, в противном случае $z_{ji} = 0$ ($i = 1, 2, 3$); результат j -го измерения — случайная величина

$$y_j = \beta_0 + \beta_1 z_{j1} + \beta_2 z_{j2} + \beta_3 z_{j3} + e_j, \quad j = 1, 2, 3, 4, \quad (1.13)$$

где β_0 — математическое ожидание показания шкалы весов при отсутствии предметов на весах, т. е.

$$\begin{aligned} \beta_0 &= E\{y_j | z_{j1} = z_{j2} = z_{j3} = 0\}; \\ Ee_j &= 0, \quad De_j = \sigma^2, \quad Ee_i e_j = 0 \quad (i \neq j). \end{aligned}$$

Для удобства заменим z_{ji} на величины x_{ji} :

$$x_{ji} = 2z_{ji} - 1, \quad i = 1, 2, 3; \quad j = 1, 2, 3, 4.$$

Получаем, что если i -й предмет в j -м измерении находился на весах, то $x_{ji} = 1$, если нет, то $x_{ji} = -1$. Положим также $x_{j0} = 1$ ($j = 1, 2, 3, 4$).

Модель (1.13) переписывается в виде

$$y_j = \theta_0 + \theta_1 x_{j1} + \theta_2 x_{j2} + \theta_3 x_{j3} + \varepsilon_j, \quad j = 1, 2, 3, 4, \quad (1.14)$$

где $\theta_0 = \beta_0 + (\beta_1 + \beta_2 + \beta_3)/2$, $\theta_i = \beta_i/2$, $i = 1, 2, 3$.

Матрица $F = \|x_{ji}\|$, $i = 0, 1, 2, 3$; $j = 1, 2, 3, 4$, определяет план эксперимента. Поскольку $|x_{ji}| = 1$ (при всех i, j), то выполняется

$$\sum_{j=1}^4 x_{ji}^2 = 4, \quad i = 0, 1, 2, 3, \quad (1.15)$$

т. е. (1.9) с $C_i^2 = 4$.

В соответствии с теоремой 1.1 минимальные дисперсии МНК-оценок $\hat{\theta}_i$ ($i = 0, 1, 2, 3$), а стало быть, и оценок $\hat{\beta}_i$ можно получить, используя любой ортогональный план, для которого $|x_{ji}| = 1$ (при всех i, j). Такой план определяется, например, матрицей

$$F = \begin{bmatrix} x_0 & x_1 & x_2 & x_3 \\ 1 & -1 & -1 & 1 \\ 1 & 1 & -1 & -1 \\ 1 & -1 & 1 & -1 \\ 1 & 1 & 1 & 1 \end{bmatrix}. \quad (1.16)$$

В силу теоремы 1.1 $D\hat{\theta}_i = \sigma^2/4$ ($i = 0, 1, 2, 3$), и, следовательно, $D\hat{\beta}_i = \sigma^2$ ($i = 1, 2, 3$). При этом

$$\hat{\beta}_1 = 2\hat{\theta}_1 = 2(-y_1 + y_2 - y_3 + y_4),$$

$$\hat{\beta}_2 = 2\hat{\theta}_2 = 2(-y_1 - y_2 + y_3 + y_4),$$

$$\hat{\beta}_3 = 2\hat{\theta}_3 = 2(y_1 - y_2 - y_3 + y_4).$$

Результаты легко обобщаются на случай произвольного числа m предметов и числа $N > m$ взвешиваний, являющегося степенью двойки (план может быть выбран как дробная реплика полного факторного эксперимента, см. гл. 2).

1.4. Непрерывные A -оптимальные планы. Теорема 1.1 следующим образом обобщается на непрерывные планы*).

Теорема 1.2. Для функции регрессии (1.1) любой ортогональный план ξ^* является A -оптимальным в множестве непрерывных планов ξ , удовлетворяющих ограничению

$$\sum_{i=1}^m \int_X x_{i\xi}^2(dx) \leq \sum_{i=1}^m \int_X x_{i\xi^*}^2(dx). \quad (1.17)$$

Доказательство аналогично доказательству теоремы 1.1. Заметим, что из утверждения теоремы не следует, что A -оптималь-

*) Непрерывные планы первого порядка определяются аналогично дискретным, см. с. 139.

ный план всегда ортогонален: в зависимости от вида множества X может оказаться, что

$$\sup_{\xi} \int_X x^2 \xi(dx)$$

в множестве всех непрерывных планов больше соответствующего супремума в множестве ортогональных планов. Лишь при совпадении указанных супремумов (такая ситуация возникает, например, когда X является m -мерным кубом, см. ниже) можно гарантировать ортогональность A -оптимального плана

1.5. D -оптимальные планы первого порядка на m -мерном кубе.

В данном пункте $X = [-1, 1]^m$, т. е. множество планирования — единичный m -мерный куб.

Сначала покажем, что непрерывным D -оптимальным является план ξ^* , сосредоточенный с равными весами во всех 2^m вершинах куба X . Эксперимент, состоящий из 2^m наблюдений и проводимый согласно этому плану, носит название *полный факторный эксперимент* (для m факторов, меняющихся на двух уровнях).

Записав матрицу (1.4) для этого плана, представленного в виде (1.2) с $N = 2^m$, видим, что выполняются (1.5) и (1.9) с $C_i^j = N$ ($i = 0, \dots, m$). Поэтому информационная матрица $M(\xi^*)$ плана ξ^* является единичной (I_{m+1}), откуда

$$d(x, \xi^*) = 1 + \sum_{i=1}^m x^2(i),$$

где $x(1), \dots, x(m)$ — координаты точки x .

Максимум

$$\max_{x \in X} d(x, \xi^*) = m + 1$$

достигается в точках плана ξ^* . На основании теоремы эквивалентности Кифера — Вольфовица (теорема 2.3 из гл. 3) заключаем, что план ξ^* является D -оптимальным. Заметим, что этот план также ортогональный, ротатабельный и A -оптимальный.

Поскольку информационные матрицы всех D -оптимальных планов совпадают, любой D -оптимальный план первого порядка на кубе ортогонален, ротатабелен и имеет единичную информационную матрицу.

Несмотря на ортогональность, ротатабельность и D -оптимальность полного факторного плана, при больших m он используется редко, так как требуемое им число 2^m измерений с ростом m растет экспоненциально, и возникают трудности, вызванные необходимостью практической реализации такого числа измерений. Для практического использования больший интерес представляют насыщенные планы, сосредоточенные в минимально возможном числе точек, равном $m + 1$. Выяснить структуру таких планов удается с помощью следующего утверждения.

Теорема 1.3. *Насыщенный план вида (1.2) первого порядка имеет информационную матрицу, пропорциональную единичной, тогда и только тогда, когда он сосредоточен в вершинах правильного m -мерного симплекса (т. е. правильного m -мерного многогранника с количеством вершин $m + 1$).*

Доказательство. Обозначим строки матрицы (1.3) через $x_{(i)}$, а строки матрицы (1.4) через $z_{(i)}$ ($i = 1, \dots, N$). По определению

$$z_{(i)} = [1 \ x_{(i)}], \quad i = 1, \dots, N,$$

а $x_{(i)}$ — это точки плана (1.2). Обозначим через φ_{jl} угол между векторами $x_{(j)}$ и $x_{(l)}$; по определению

$$\cos \varphi_{jl} = \frac{(x_{(j)}, x_{(l)})}{\|x_{(j)}\| \cdot \|x_{(l)}\|}, \quad \text{где } (x_{(j)}, x_{(l)}) = \sum_{i=1}^m x_{ji}x_{li}, \quad \|x\| = \sqrt{(x, x)}.$$

То, что информационная матрица плана (1.2) пропорциональна единичной, записывается в виде

$$(z_{(j)}, z_{(l)}) = \sum_{i=0}^m x_{ji}x_{li} = \begin{cases} 0, & \text{если } j \neq l, \\ C_1, & \text{если } j = l, \end{cases} \quad (1.18)$$

где $C_1 > 1$ — некоторая константа (если информационная матрица равна единичной, то $C_1 = N$). Поскольку $x_{j0} = 1$ ($j = 1, \dots, N$), то (1.18) эквивалентно тому, что

$$(x_{(j)}, x_{(l)}) = \begin{cases} -1, & \text{если } j \neq l, \\ C_1 - 1, & \text{если } j = l, \end{cases}$$

т. е. тому, что

$$\begin{aligned} \|x_{(l)}\| &= \sqrt{(x_{(l)}, x_{(l)})} = \sqrt{C_1 - 1}, \\ \cos \varphi_{jl} &= -1/(C_1 - 1), \quad j \neq l, \quad j, l = 1, \dots, N. \end{aligned}$$

Это эквивалентно тому, что точки $x_{(1)}, \dots, x_{(N)}$ лежат в вершинах правильного m -мерного симплекса. Теорема доказана.

Планы, сосредоточенные в вершинах правильного симплекса, называются *симплекс-планами*.

Из теоремы 1.3 следует, что D -оптимальный насыщенный план первого порядка на кубе существует, если существует правильный симплекс, все вершины которого совпадают с некоторыми из вершин куба. Однако этого можно достичь не для всех размерностей. Например, при $m = 3$ это возможно, а при $m = 2$ — нет.

Построение правильных симплексов с указанным свойством эквивалентно построению $(m + 1) \times (m + 1)$ -матриц Адамара. Приведем некоторые сведения об этих матрицах.

Матрицей Адамара называется квадратная $n \times n$ -матрица $A = \|a_{ij}\|$, элементы a_{ij} которой равны $+1$ или -1 и выполнено равенство

$$1 \cdot 1^T = nI_n \quad (1.19)$$

Свойства матриц Адамара:

1) матрица Адамара имеет максимальное по модулю значение определителя в множестве всевозможных квадратных матриц того же порядка с элементами, по модулю не превосходящими единицы;

2) любые две строки матрицы Адамара ортогональны;

3) из $AA^T = nI_n$ следует $A^T A = nI_n$, и наоборот;

4) перестановка строк или столбцов и их умножение на -1 приводит к матрице Адамара;

5) если $A = \|a_{ij}\|$ и B — матрицы Адамара порядков $m \times m$ и $n \times n$ соответственно, то матрица $C = \|a_{ij}B\|$ является матрицей Адамара порядка $mn \times mn$.

Матрица Адамара, у которой первая строка и первый столбец состоят из $+1$, называется *нормализованной*. Именно такие матрицы порядка $(m+1) \times (m+1)$ и могут быть выбраны в качестве матрицы плана F . Матрицы Адамара порядка $n \times n$ построены, например, для всех n вида $n = 4k$ ($k = 1, \dots, 50$, кроме $k = 47$). Методы построения матриц Адамара содержатся в [43].

1.6. Общее свойство непрерывных D -оптимальных планов.

Теорема 1.4. Пусть множество планирования X — компактное подмножество R^n . Тогда все точки D -оптимального плана первого порядка лежат на границе множества X .

Доказательство. Обозначим через $x(i)$ координаты точки $x \in X$ ($i = 1, \dots, m$) и $x(0) = 1$. Тогда базисные функции $f_i(x)$ равны $f_i(x) = x(i)$.

Пусть ξ — произвольный план первого порядка. Тогда

$$\begin{aligned} d(x, \xi) &= f^T(x) M^{-1}(\xi) f(x) = \\ &= (1, x(1), \dots, x(m)) M^{-1}(\xi) (1, x(1), \dots, x(m))^T = \\ &= \sum_{i=0}^m (x(i))^2 d_{ii}(\xi) + \sum_{\substack{i, j=0 \\ i \neq j}}^m x(i) x(j) d_{ij}(\xi), \end{aligned}$$

где $d_{ij}(\xi)$ — элементы матрицы $D(\xi)$. Отсюда

$$\frac{\partial^2 d(x, \xi)}{\partial (x(i))^2} = 2d_{ii}(\xi) > 0.$$

и, следовательно, максимум функции $d(x, \xi)$ не может достигаться ни в одной внутренней точке множества X . Но по теореме эквивалентности Кифера — Вольфовица (теорема 2.3 из гл. 3) для D -оптимального плана ξ^* максимум функции $d(x, \xi^*)$ достигается в точках этого плана. Поэтому все точки плана ξ^* должны лежать на границе множества X . Теорема доказана.

Упражнения.

1. Пусть дана схема регрессии (1.13). Сравните точность МНК-оценок, получаемых при использовании планов, определяемых матрицами (1.16) и

$$\begin{bmatrix} 1 & 1 & 1 & -1 \\ 1 & -1 & 1 & 1 \\ 1 & 1 & -1 & 1 \\ 1 & 1 & 1 & 1 \end{bmatrix}, \quad \begin{bmatrix} 1 & 1 & -1 & 1 \\ 1 & -1 & 1 & -1 \\ 1 & 1 & 1 & -1 \\ 1 & -1 & -1 & 1 \end{bmatrix}, \quad \begin{bmatrix} 1 & -1 & -1 & 1 \\ 1 & 1 & -1 & -1 \\ 1 & -1 & 1 & -1 \\ 1 & -1 & -1 & -1 \end{bmatrix}.$$

2. Докажите, что план, определяемый матрицей (1.16), является полуреplikой факторного эксперимента 2^3 , и выпишите для этой полуреплики генерирующее и определяющее соотношения.

3. Постройте матрицы Адамара порядков 2×2 и 4×4 .
4. Приведите пример A -оптимального плана, который не является ортогональным для функции регрессии (1.1).
5. Докажите теорему 1.2.
6. Докажите, что симплекс-планы являются не только D -оптимальными, но и A -оптимальными

§ 2. Некоторые обобщения классической постановки задачи планирования регрессионного эксперимента

2.1. Планирование эксперимента с областью действия в функциональном пространстве. Классическая теория планирования регрессионного эксперимента (см. гл. 3) развита в рамках схемы регрессии, в которой аргумент x линейной по параметрам функции регрессии $\eta(x, \theta)$ выбирается из некоторого множества планирования X , а $\theta = (\theta_1, \dots, \theta_m)^T$ — вектор неизвестных параметров. Формально множество планирования X не обязано быть конечномерным, однако большинство конкретных результатов классической теории (включая численные методы построения планов и практические приложения) относится к случаю, когда это множество конечномерно.

Для широкого круга задач планирования, связанных с постановкой физических экспериментов, типичной является ситуация, когда единичное наблюдение реализуется с помощью функционала $x: \{f\} \rightarrow \mathbb{R}$, сопоставляющего состоянию наблюдаемого объекта со значениями некоторой вещественной переменной, т. е. отображающего множество состояний $\{f\}$ в вещественную прямую. В тех случаях, когда состояние объекта наблюдения описывается некоторой функцией $f(r, t, \dots)$ пространственных, временных и т. п. координат и функция f принадлежит определенному функциональному классу $\mathcal{F}$, указанное отображение задает соответствующий функционал над классом $\mathcal{F}$. Следовательно, в описанной ситуации множество планирования X является некоторым множеством функционалов на функциональном классе $\mathcal{F}$.

Метрика на множестве планирования X индуцируется метрикой исходного функционального класса $\mathcal{F}$, согласованной с физической природой задачи.

Если $\mathcal{F}$ есть полное нормированное (банахово) пространство (что далее и предполагается), то, ограничиваясь линейными функционалами, получаем в качестве множества планирования сопряженное банахово пространство $\mathcal{F}^*$ (пространство ограниченных непрерывных линейных функционалов на $\mathcal{F}$). Формулировка теории планирования эксперимента на языке пары сопряженных функциональных пространств $(\mathcal{F}, \mathcal{F}^*)$ является наиболее естественной при изучении бесконечномерных задач планирования эксперимента средствами функционального анализа.

В терминах сопряженных функциональных пространств $(\mathcal{F}, \mathcal{F}^*)$ задача планирования регрессионного эксперимента, со

ответствующая классической схеме линейной по параметрам регрессии, формулируется следующим образом. Пусть $\mathcal{F}$ — полное нормированное пространство и $L \subset \mathcal{F}$ — его m -мерное линейное подпространство. Выбирая в L базис, т. е. набор линейно независимых элементов $\{e_1, \dots, e_m\}$, каждый элемент $f \in L$ можно представить в виде линейной комбинации

$$f = \sum_{i=1}^m \theta_i e_i. \quad (2.1)$$

Эксперимент, предназначенный для оценивания параметров $\theta = (\theta_1, \dots, \theta_m)^T$, состоит в наблюдении значений случайных величин $y_1, \dots, y_N$:

$$y_j = x_j(f) + e_j, \quad j = 1, \dots, N, \quad (2.2)$$

где x_j — линейные ограниченные функционалы на $\mathcal{F}$, т. е. элементы сопряженного банахова пространства $\mathcal{F}^*$, а случайные ошибки e_j , как обычно, центрированы, некоррелированы и имеют одинаковые конечные дисперсии. С точки зрения оценивания параметров $\theta_1, \dots, \theta_m$ задача (2.1), (2.2) эквивалентна обычной (конечномерной) задаче оценивания параметров линейной регрессии

$$y_j = \sum_{i=1}^m \theta_i x_j(e_i) + e_j, \quad j = 1, \dots, N, \quad (2.3)$$

с МНК-оценками в качестве наилучших линейных несмещенных оценок.

Для того чтобы привести (2.3) к рассматривавшейся в гл. 1 классической линейной регрессионной модели, достаточно положить $x_j = x_j(e_i)$ ($j = 1, \dots, N$; $i = 1, \dots, m$). Аналогично гл. 3 непрерывным планом эксперимента ξ будем называть набор

$$\xi = \left\{ \begin{matrix} x_1, \dots, x_n \\ p_1, \dots, p_n \end{matrix} \right\} \quad (2.4)$$

функционалов x_j и их весов p_j :

$$\sum_{j=1}^n p_j = 1, \quad p_j \geq 0, \quad x_j \in X, \quad j = 1, \dots, n.$$

Множеством планирования является подмножество X сопряженного банахова пространства $\mathcal{F}^*$, которому по условию принадлежат функционалы x_j плана ξ .

План эксперимента можно рассматривать и как вероятностную меру, сосредоточенную не обязательно на конечном подмножестве множества X .

Информационной матрицей плана (2.4) является матрица $M(\xi)$ с элементами

$$M_{ik}(\xi) = \sum_{j=1}^n p_j x_j(e_i) x_j(e_k), \quad i, k = 1, \dots, n. \quad (2.5)$$

Если множество планирования X ограничено и замкнуто в норме сопряженного пространства $\mathcal{F}^*$, то семейство матриц

$\mathcal{M} = \{M(\xi)\}$ образует выпуклое компактное множество и справедливы остальные утверждения гл. 3 относительно информационных матриц. В частности, для любого плана ξ элементы матрицы $M(\xi)$ могут быть представлены в виде (2.5) с $n \leq m(m+1)/2 + 1$.

Как и в гл. 3, оптимизация плана ξ может проводиться как на основе характеристик точности НЛН-оценок параметров θ , так и с точки зрения точности оценивания функционалов из некоторого множества $Y \subset \mathcal{F}^*$ (указанное множество не обязательно, вообще говоря, совпадает с X). В соответствии с этим получаем аналоги основных критериев оптимальности планов. Так, D -оптимальным является план

$$\xi^* = \arg \max_{\xi} \det M(\xi).$$

В случае невырожденности матрицы $M(\xi)$ для любого функционала $x \in \mathcal{F}^*$ числовая функция

$$d(x, \xi) = \sum_{i, k=1}^m [M^{-1}(\xi)]_{ik} x(e_i) x(e_k)$$

определяет дисперсию НЛН-оценки функционала x по результатам эксперимента (2.3), а план

$$\xi^* = \arg \min_{\xi} \{ \max_{x \in X} d(x, \xi) \}$$

является G -оптимальным.

В определенном смысле описанная формулировка задачи планирования экспериментов шире классической, так как для получения последней требуется в качестве пространства $\mathcal{F}$ выбрать пространство C_U функций $f(u)$, непрерывных на некотором компакте U , а в качестве «функциональной» области планирования $X \subset C_U^*$ — множество функционалов вида $x_u(f) = f(u)$, $u \in U$, где U — «обычная» область планирования (например, ограниченное множество в пространстве $\mathbb{R}^m$).

С точки зрения оценивания параметров θ в схеме измерений (2.3) расширение множества планирования X до подмножества сопряженного функционального пространства не дает ничего нового — каждая точка (т. е. функционал $x \in X$) представлена в схеме измерений и в информационной матрице плана $M(\xi)$ только своими значениями на элементах базиса $\{e_1, \dots, e_n\}$. Покажем, что и задача планирования в определенном смысле эквивалентна конечномерной.

Фиксируем базис подпространства $L \subset \mathcal{F}$ в разложении (2.1) и рассмотрим отображение сопряженного пространства $\mathcal{F}^*$ в m -мерное евклидово пространство $\mathbb{R}^m$:

$$\begin{aligned} \varphi: \mathcal{F}^* &\rightarrow \mathbb{R}^m: x \rightarrow \varphi(x) = (z_1, \dots, z_m)^T, \\ z_i &= x(e_i), \quad i = 1, \dots, m. \end{aligned} \quad (2.6)$$

Отображение (2.6) сопоставляет каждому функционалу $x \in \mathcal{F}^*$ вектор $\varphi(x)$, составленный из значений этого функционала на элементах базиса $\{e_1, \dots, e_m\}$. Отображение (2.6) непрерывно и переводит каждое ограниченное замкнутое множество в компактное. В частности, ограниченное замкнутое множество планирования X отображается в компакт $Z = \{\varphi(x), x \in X\} \subset \mathbb{R}^m$. При этом информационные матрицы оптимальных по любому критерию планов для множества планирования $X \subset \mathcal{F}^*$ в функциональном пространстве совпадают с информационными матрицами соответствующих планов для линейной регрессии вида

$$Ey(z) = \sum_{i=1}^m \theta_i z_i \quad (2.7)$$

на множестве планирования $Z = \varphi(X) \subset \mathbb{R}^m$.

Задача оптимального планирования на компактном множестве $Z \subset \mathbb{R}^m$ для функции регрессии (2.7) может быть решена с помощью теоретических результатов и численных методов, содержащихся в гл. 3. Отметим, что при решении задачи оптимального планирования могут встретиться дополнительные трудности, поскольку сложным может быть вид множества $Z = \varphi(X)$.

За исключением простейших случаев, указанный подход еще не обеспечивает окончательного решения задачи оптимального планирования. Действительно, с помощью отображения (2.6) бесконечномерная задача планирования сводится к конечномерной в том смысле, что решение последней позволяет найти информационную матрицу оптимального «бесконечномерного» плана. Однако в конкретных задачах этого недостаточно: необходимо реализовать оптимальный план с помощью функционалов — элементов функционального пространства $\mathcal{F}^*$.

Пусть в результате решения конечномерной задачи для области $Z = \varphi(X) \subset \mathbb{R}^m$ получен оптимальный план ξ_z^* . Как известно (см. с. 97), этот план всегда можно считать сосредоточенным на конечном множестве точек, число которых не превышает $m(m+1)/2$. Пусть $z^* = (z_1^*, \dots, z_m^*) \in Z$ — одна из таких точек с весом $p^* > 0$; тогда для восстановления соответствующей ей «функциональной» точки необходимо найти функционал, удовлетворяющий системе уравнений

$$x^*(e_i) = z_i^*, \quad i = 1, \dots, m, \quad (2.8)$$

и принадлежащий множеству $X \subset \mathcal{F}^*$. Приписав такому функционалу x^* вес p^* и повторив эту операцию для каждой точки, входящей в план ξ_z^* с ненулевым весом, получим план в X с той же информационной матрицей и поэтому оптимальный в смысле того же критерия, что и план ξ_z^* .

В различных приложениях естественным является различный выбор пары пространств $(\mathcal{F}, \mathcal{F}^*)$. Наиболее часто в качестве $\mathcal{F}$ используют L_2 — пространство функций, суммируемых

с квадратом. С математической точки зрения этот случай прост, поскольку $\mathcal{F} = \mathcal{F}^*$. На нем мы далее и остановимся.

Будем предполагать, что $\mathcal{F} = L_2(T, \mu)$ — гильбертово пространство функций, интегрируемых с квадратом относительно некоторой меры μ на компакте T . Норма функции $f \in \mathcal{F}$ определяется по формуле

$$\|f\|^2 = \int_T |f(t)|^2 \mu(dt).$$

В этом случае $\mathcal{F}^*$ — тоже пространство $L_2(T, \mu)$ и для любых $x \in \mathcal{F}^*$, $f \in \mathcal{F}$ выполнено равенство

$$x(f) = \int_T x(t) f(t) \mu(dt), \quad \|x\|^2 = \int_T |x(t)|^2 \mu(dt).$$

Пример 1.1. Пусть множество планирования X — единичный шар в сопряженном пространстве $\mathcal{F}^* = L_2(T, \mu)$:

$$X = \{x \in L_2(T, \mu) \mid \|x\| < 1\}.$$

Функции $e_1, \dots, e_m$ ортонормированы в $L_2(T, \mu)$:

$$\int_T e_i(t) e_k(t) \mu(dt) = \begin{cases} 1, & i = k, \\ 0, & i \neq k. \end{cases}$$

В качестве конечномерного множества Z получаем

$$Z = \{z = (z_1, \dots, z_m)^T \mid z_1^2 + \dots + z_m^2 \leq 1\},$$

т. е. единичный шар пространства $\mathbb{R}^m$.

Как известно (см. § 1), оптимальный план ξ_z^* для линейной регрессии на шаре сосредоточен с равными весами в $m+1$ точках, лежащих на поверхности шара в вершинах правильного симплекса. Пусть $z^* = (z_1^*, \dots, z_m^*)^T$ — одна из таких точек. Этой точке соответствует функционал x^* , определяемый функцией

$$x^*(t) = \sum_{i=1}^m z_i^* e_i(t).$$

Для этого функционала справедливость системы уравнений (2.8) очевидна.

Построив указанным образом $m+1$ функций $x_i^*(t)$, соответствующих всем вершинам правильного симплекса, и приписав им равные веса $(m+1)^{-1}$, получим D -оптимальный план.

Результат примера может быть существенно обобщен: для случая $X \subset L_2(T, \mu)$ бесконечномерная задача планирования всегда легко сводится к конечномерной.

Пусть $x \in L_2(T, \mu)$, L — m -мерное подпространство пространства $L_2(T, \mu)$, $\{e_1, \dots, e_m\}$ — μ -ортонормированный базис под-

пространства L . Тогда проекция $x(L)$ точки x на L записывается в виде

$$x(L) = \sum_{i=1}^m e_i x(e_i).$$

Теорема 2.1. Пусть $X \subset L_2(T, \mu)$. Тогда информационная матрица $M(\xi)$ любого плана ξ вида (2.4) совпадает с информационной матрицей $M(\xi(L))$ плана

$$\xi(L) = \begin{pmatrix} x_1(L), \dots, x_n(L) \\ p_1, \dots, p_n \end{pmatrix},$$

где

$$x_j(L) = \sum_{i=1}^m e_i x_j(e_i), \quad j = 1, \dots, n,$$

суть проекции точек $x_j \in L_2(T, \mu)$ на m -мерное пространство, порожденное ортонормированным базисом $\{e_1, \dots, e_m\}$.

Доказательство. Поскольку $x_j(e_i)$ — проекция точки x_j на элемент базиса e_i , то для всех $i = 1, \dots, m$, $j = 1, \dots, n$ выполнено равенство

$$x_j(e_i) = [x_j(L)](e_i). \quad (2.9)$$

По определению информационная матрица $M(\xi(L))$ плана $\xi(L)$ состоит из чисел

$$M_{ik}(\xi(L)) = \sum_{j=1}^n p_j [x_j(L)](e_i) [x_j(L)](e_k),$$

где $i, k = 1, \dots, m$. С учетом (2.9) эти числа совпадают с (2.5) — элементами информационной матрицы $M(\xi)$. Теорема доказана.

Из доказанной теоремы вытекает, что без потери информации (т. е. без изменения множества информационных матриц) в случае $X \subset L_2(T, \mu)$ можно ограничиваться планами эксперимента, сосредоточенными на конечномерном множестве $X \cap L$.

2.2. Планирование регрессионного эксперимента при коррелированных наблюдениях. Следующий класс задач, имеющих важное прикладное значение, составляют задачи с коррелированными ошибками наблюдений. В этих задачах результаты измерений имеют аддитивную погрешность, представляющую собой значение реализации случайного процесса в соответствующей точке.

Сначала остановимся на случае, когда имеется только один неизвестный параметр. Предположим, что на промежутке $T = [0, 1]$ задан случайный процесс

$$y(t) = \theta f(t) + e(t), \quad t \in T, \quad (2.10)$$

где $f(t)$ — известная непрерывная функция на T ; θ — скалярный неизвестный параметр; $e(t)$ — случайный процесс с нулевым средним, известной ковариационной функцией

$$K(s, t) = E[e(s)e(t)] \quad (2.11)$$

и траекториями, непрерывными с вероятностью единица. Обозначим через T_N произвольное множество из N точек промежутка T :

$$T_N = \{t_1, \dots, t_N | 0 \leq t_1 < t < \dots < t_N \leq 1\}. \quad (2.12)$$

Будем предполагать, что функция $K(s, t)$ образует невырожденную матрицу при замене множества T на произвольное множество вида (2.12), т. е. что $\det K_N \neq 0$, где $K_N = \|K(t_i, t_k)\|_{i, k=1}^N$, $t_i \in T_N$ ($i = 1, \dots, N$). Пусть множество T_N фиксировано и определяет точки проведения измерений, т. е. план эксперимента для N измерений. Тогда вектором результатов измерений является вектор

$$Y_N = (y(t_1), \dots, y(t_N))^T.$$

Ковариационной матрицей этого вектора является матрица K_N .

Таким образом, в обозначениях гл. 1 имеем линейную регрессионную модель $(Y_N, F_N \theta, K_N)$, где $F_N = (f(t_1), \dots, f(t_N))^T$. Согласно результатам п. 1.7 гл. 1, НЛН-оценкой параметра θ является

$$\hat{\theta}_N = (F_N^T K_N^{-1} F_N)^{-1} F_N^T K_N^{-1} Y_N. \quad (2.13)$$

Дисперсия этой оценки равна

$$D \hat{\theta}_N = (F_N^T K_N^{-1} F_N)^{-1}. \quad (2.14)$$

Примем обозначение

$$\|f\|^2 = \sup_{N, T_N} F_N^T K_N^{-1} F_N,$$

где супремум берется по всевозможным конечным набором T_N вида (2.12).

Определим в множестве ограниченных по указанной норме функций $\mathcal{F} = \{f | \|f\| < \infty\}$ скалярное произведение по формуле

$$(f, g)_K = \sup_{T_N} F_N^T K_N^{-1} G_N,$$

где $G_N = (g(t_1), \dots, g(t_N))^T$, $t_i \in T_N$ ($i = 1, \dots, N$). Легко проверить (проверьте!), что множество функций $\mathcal{F}$ удовлетворяет следующим двум свойствам: а) при любом $t \in T$ функция $K(\cdot, t)$ принадлежит $\mathcal{F}$; б) для любой функции $f \in \mathcal{F}$ и любого $t \in T$ выполнено равенство

$$(f, K(\cdot, t))_K = f(t). \quad (2.15)$$

Множество функций с указанными свойствами называется *гильбертовым пространством с воспроизводящим ядром*, а свойство б) называется *воспроизводящим свойством ядра K* .

Из свойств гильбертовых пространств с воспроизводящим ядром вытекает, что пространство $\mathcal{F}$ сепарабельно, содержит только непрерывные функции и порождается всевозможными линейными комбинациями $\sum a_i K(s_i, t)$, $t \in T$.

Рассмотрим проблему оптимального выбора плана эксперимента, т. е. конечного набора T_N . Пусть $S_N = \{T_N\}$ — множество всех конечных наборов вида (2.12), содержащих ровно N различных точек. Оптимальным дискретным планом для модели (2.10) является план

$$T_N^* = \underset{T_N \in S_N}{\text{Arg sup}} \|f\|_{T_N}, \quad (2.16)$$

где норма $\|\cdot\|_{T_N}$, соответствующая плану T_N , порождается квадратичной формой с матрицей K_N^{-1} :

$$\|f\|_{T_N}^2 = F_N^T K_N^{-1} F_N. \quad (2.17)$$

Из (2.14) следует, что для плана (2.16) дисперсия НЛН-оценки неизвестного параметра θ минимальна в множестве всех планов вида (2.12) при фиксированном N .

Дискретные оптимальные планы (2.16) могут быть найдены путем прямой максимизации критерия (2.17) по точкам $t_1, \dots, t_N$ плана (2.12). При этом типичной является такая ситуация, когда оптимальный план T_{N+1}^* для проведения $(N+1)$ -го измерения лучше оптимального плана T_N^* для N измерений, т. е.

$$\|f\|_{T_{N+1}^*} > \|f\|_{T_N^*}.$$

Исключением является случай, когда при некотором N функция f из (2.10) представима в виде

$$f(t) = \sum_{i=1}^N a_i K(t, t_i). \quad (2.18)$$

В этом случае

$$\|f\| = \|f\|_{T_N^*}, \quad (2.19)$$

причем оптимальный план T_N^* состоит из точек t_i , фигурирующих в (2.18). План T_N^* с наименьшим возможным N , для которого имеет место (2.19), называется *глобально оптимальным*. Очевидно, что если ни при каком конечном N представления (2.18) не существует, то глобально оптимального плана также не существует.

Поскольку для всех $f \in \mathcal{F}$ справедливо равенство

$$\|f\| = \lim_{N \rightarrow \infty} \sup_{T_N \in S_N} \|f\|_{T_N}$$

то представляется естественным ослабить требование точной оптимальности плана при каждом N до асимптотической оптимальности последовательности планов $\{T_N, N \rightarrow \infty\}$. Пусть в модели (2.10) $f(t)$ — непрерывная функция, допускающая представление

$$f(t) = \int_0^1 K(s, t) \varphi(s) ds, \quad (2.20)$$

где функция $\varphi(s)$ также непрерывна на $T = [0, 1]$ (поэтому, в частности, $f \in \mathcal{F}$). Последовательность планов $\{T_N, N \geq 1\}$ называется *асимптотически оптимальной*, если

$$\lim_{N \rightarrow \infty} \frac{\|f\|^2 - \inf_{T_N} \|f\|_{T_N}^2}{\sup_{T'_N \in S_N} \|f\|_{T'_N}^2} = 1.$$

Без доказательства приведем теорему, результат которой описывает структуру асимптотически оптимальной последовательности планов.

Теорема 2.2. Пусть ядро $K(s, t)$ непрерывно на квадрате $T \times T$ и имеет непрерывные производные до второго порядка включительно во всех точках квадрата вне главной диагонали ($s \neq t$); на диагонали $s = t$ функция $K(s, t)$ имеет все правые и левые производные до второго порядка включительно и ненулевой скачок первой производной: функция

$$\alpha(t) = \lim_{s \downarrow t} \frac{\partial}{\partial s} K(s, t) - \lim_{s \uparrow t} \frac{\partial}{\partial s} K(s, t) \quad (2.21)$$

строго положительна и непрерывна на T . Пусть далее $\partial K(\cdot, t)/\partial t^2 \in \mathcal{F}$ при любом $t \in T$ и нормы этого семейства функций ограничены в совокупности.

Тогда асимптотически оптимальная последовательность $\{T_N, N \rightarrow \infty\}$ определяется через функцию $h(t) = [\alpha(t)\varphi^2(t)]^{1/3}$ следующим образом:

$$\int_0^{t_i} h(t) dt = \frac{i-1}{N-1} \int_0^1 h(t) dt, \quad i = 1, \dots, N,$$

причем t_i^* — наименьшее число, удовлетворяющее написанному условию.

Приведенные в формулировке теоремы условия на ядро $K(s, t)$ связаны с разрешимостью интегрального уравнения (2.20) относительно функции $\varphi(s)$ в множестве непрерывных функций. Примером семейства ядер, удовлетворяющих указанным условиям, является

$$K(s, t) = \int_0^\infty \exp\{-x|t-s|\} p(x) dx, \quad (2.22)$$

где $p(x)$ — плотность распределения вероятностей на $[0, \infty)$ с конечным третьим моментом.

Ряд приведенных выше результатов относительно модели (2.10) может быть обобщен на случай, когда имеется несколько неизвестных параметров, т. е. вместо (2.10) справедлива модель

$$y(t) = \theta^T f(x) + \varepsilon(t), \quad t \in T = [0, 1],$$

где $\theta = (\theta_1, \dots, \theta_m)^T$ — вектор неизвестных параметров, $f(x) = (f_1(x), \dots, f_m(x))^T$ — вектор базисных функций, а предположения относительно случайного процесса $\varepsilon(t)$ те же, что и раньше. Явный вид НЛН-оценок $\hat{\theta}_N$ параметров θ определяется, так же как и раньше, по формуле (2.13), а дисперсионная матрица $D\hat{\theta}_N$ оценок $\hat{\theta}_N$ — по формуле (2.14).

В отличие от случая скалярного параметра θ выражение (2.14) является не числом, а матрицей, и задача ее минимизации является неопределенной. Можно лишь, так же как и в гл. 3,

минимизировать выпуклые функционалы на множестве дисперсионных матриц. Для того чтобы соответствующие критерии качества имели смысл, необходимо, как и в одномерном случае, чтобы все функции $f_i(t)$ ($i = 1, \dots, m$) принадлежали гильбертову пространству с воспроизводящим ядром K .

Относительно дискретных оптимальных планов в этом случае, так же как и для случая скалярного параметра, ничего не известно, за исключением ситуации, когда существует глобально оптимальный план, т. е. когда все функции $f_i(t)$ представимы в виде конечных линейных комбинаций $\{K(t_j, t), t_j \in T_N\}$. Асимптотически оптимальные последовательности планов определяются, как в одномерном случае, только вместо функционалов типа $\|f\|_{H_N}$ используется некоторый критерий $\Psi(D\hat{\theta}_N)$. Для построения таких последовательностей также используется представление через квантили некоторой положительной функции $h(t)$, определяемой по ядру $K(s, t)$, по совокупности решений $\varphi_i(s)$ интегральных уравнений (2.20) для всех f_i и по критерию Ψ .

Рассмотренные выше подходы к решению задачи оптимального планирования эксперимента по оцениванию параметров случайных процессов существенно отличаются от классического подхода, рассмотренного в гл. 3. Для некоторых достаточно простых постановок могут быть получены результаты, похожие на результаты классической теории. Ниже рассмотрена одна из таких постановок.

Рассмотрим задачу оценивания неизвестного среднего θ случайного процесса

$$y(t) = \theta + \varepsilon(t), \quad t \in T = [0, 1], \quad (2.23)$$

где, как и выше, $\varepsilon(t)$ — случайный процесс с нулевым средним и ковариационной функцией (2.11), а множеством планирования X является $T = [0, 1]$. Модель (2.23) является частным случаем модели (2.10), получающимся при $f(t) = 1$. Для оценивания параметра θ будем использовать не НЛН-оценку, как выше, а стандартную МНК-оценку, которая по результатам измерений $y(t_i)$ в точках $t_i \in T_N$ из набора (2.12) строится очень просто:

$$\hat{\theta}_N = \frac{1}{N} \sum_{i=1}^N y(t_i), \quad (2.24)$$

т. е. $\hat{\theta}_N$ является средним арифметическим результатов измерений. Дисперсия этой оценки равна

$$\begin{aligned} D\hat{\theta}_N &= E \left[\frac{1}{N} \sum_{j=1}^N \varepsilon(t_j) \right]^2 = \frac{1}{N^2} \sum_{i,j=1}^N E \varepsilon(t_i) \varepsilon(t_j) = \\ &= \frac{1}{N^2} \sum_{i,j=1}^N K(t_i, t_j). \end{aligned} \quad (2.25)$$

Аналогично классической теории планирования регрессионного эксперимента введем в рассмотрение непрерывные планы эксперимента на $X = T$, т. е. произвольные вероятностные меры ξ на борелевских подмножествах множества $T = [0, 1]$. Множество непрерывных планов обозначим, как и в гл. 3, через Ξ .

Если план эксперимента имеет вид

$$\xi_N = \left\{ \frac{t_1}{1/N}, \dots, \frac{t_N}{1/N} \right\}, \quad (2.26)$$

то измерения случайного процесса $y(t)$ проводятся в точках $t_1, \dots, t_N$ (т. е. в точках множества (2.12)). Положим

$$D(\xi) = \iint_{T \times T} K(s, t) \xi(ds) \xi(dt). \quad (2.27)$$

Для случая (2.26) выражения (2.25) и (2.27) совпадают, и, следовательно, $D(\xi)$ можно рассматривать как дисперсию оценки $\int y(t) \xi(dt)$ (непрерывного аналога оценки (2.24)) и как критерий качества плана эксперимента ξ , который необходимо минимизировать на множестве Ξ .

Нашей основной целью является доказательство теоремы типа теоремы эквивалентности. Для этого необходимо, во-первых, доказать, что функционал $D(\xi)$ является выпуклым на множестве непрерывных планов Ξ , и, во-вторых, уметь вычислять производные этого функционала по направлениям. Этим вспомогательным целям служат две приводимые ниже леммы.

Лемма 2.1. *Функционал (2.27) является выпуклым на множестве непрерывных планов Ξ , т. е. для плана*

$$\xi_\alpha = (1 - \alpha) \xi_0 + \alpha \xi_1 \quad (2.28)$$

выполнено неравенство

$$D(\xi_\alpha) \leq (1 - \alpha) D(\xi_0) + \alpha D(\xi_1), \quad (2.29)$$

где ξ_0, ξ_1 — произвольные непрерывные планы, а α — любое число из промежутка $[0, 1]$.

Доказательство. Используя симметричность функции $K(s, t)$ (т. е. то, что $K(s, t) = K(t, s)$ для всех $s, t \in T$), имеем

$$\begin{aligned} D(\xi_\alpha) &= \iint K(s, t) [(1 - \alpha) \xi_0(ds) + \alpha \xi_1(ds)] [(1 - \alpha) \xi_0(dt) + \\ &\quad + \alpha \xi_1(dt)] = (1 - \alpha)^2 \iint K(s, t) \xi_0(ds) \xi_0(dt) + \\ &\quad + 2\alpha(1 - \alpha) \iint K(s, t) \xi_0(ds) \xi_1(dt) + \\ &\quad + \alpha^2 \iint K(s, t) \xi_1(ds) \xi_1(dt) = (1 - \alpha) D(\xi_0) + \alpha D(\xi_1) - \alpha(1 - \alpha) A, \end{aligned}$$

где

$$A = \iint K(s, t) [\xi_0(ds) \xi_0(dt) + \xi_1(ds) \xi_1(dt) - 2\xi_0(ds) \xi_1(dt)].$$

Для доказательства справедливости неравенства (2.29) достаточно показать, что $A \geq 0$. Для этого обозначим через $p_0(s)$ и $p_1(s)$ плотности вероятностных мер $\xi_0(ds)$ и $\xi_1(ds)$ относительно некоторой вероятностной меры $\nu(ds)$ (в качестве ν можно взять, например, $\nu = (\xi_0 + \xi_1)/2$), после чего преобразуем выражение для A :

$$A = \iint K(s, t) [p_0(s)p_0(t) + p_1(s)p_1(t) - 2p_0(s)p_1(t)] \nu(ds)\nu(dt) = \\ = \iint K(s, t) a(s)a(t) \nu(ds)\nu(dt),$$

где $a(s) = p_0(s) - p_1(s)$. Это выражение неотрицательно в силу свойств ковариационной функции $K(s, t)$. Действительно, если мера ν сосредоточена на конечном числе точек $t_1, \dots, t_N$, то неотрицательность A вытекает из неотрицательной определенности матрицы $\|K(t_i, t_j)\|_{i,j=1}^N$. В общем случае A представляет собой дисперсию непрерывной МНК-оценки $\int z(t) \nu(dt)$ параметра θ , полученной по плану $\nu(dt)$ для процесса $z(t) = \theta + a(t)e(t)$. Лемма доказана.

Лемма 2.2. Пусть план ξ_α имеет вид (2.28). Тогда

$$\frac{\partial D(\xi_\alpha)}{\partial \alpha} \Big|_{\alpha=0} = 2 \left[\iint K(s, t) \xi_0(ds) \xi_1(dt) - D(\xi_0) \right]. \quad (2.30)$$

Доказательство. Используем полученное при доказательстве леммы 2.1 выражение для $D(\xi_\alpha)$. Имеем

$$\frac{\partial D(\xi_\alpha)}{\partial \alpha} \Big|_{\alpha=0} = \frac{\partial}{\partial \alpha} \left[(1-\alpha)^2 D(\xi_0) + \alpha^2 D(\xi_1) + \right. \\ \left. + 2\alpha(1-\alpha) \iint K(s, t) \xi_0(ds) \xi_1(dt) \right] \Big|_{\alpha=0} = \\ = -2D(\xi_0) + 2 \iint K(s, t) \xi_0(ds) \xi_1(dt).$$

Лемма доказана.

Используем теперь необходимое и достаточное условие оптимальности для выпуклого функционала, аналогичное условию (2.1) из гл. 3. В данном случае план ξ^* оптимален, если и только если для любого плана ξ выполняется неравенство

$$\frac{\partial D((1-\alpha)\xi^* + \alpha\xi)}{\partial \alpha} \Big|_{\alpha=0} \geq 0. \quad (2.31)$$

Учитывая (2.30), перепишем (2.31) в виде

$$\inf_{\xi \in \mathcal{Z}} \iint K(s, t) \xi^*(ds) \xi_1(dt) \geq D(\xi^*). \quad (2.32)$$

Минимум в левой части (2.32) достигается на плане ξ_1 , сосредоточенном в точке минимума функции

$$\varphi(t) = \int K(s, t) \xi^*(ds),$$

и поэтому необходимое и достаточное условие оптимальности (2.32) переписывается в виде

$$\inf_{t \in T} \int K(s, t) \xi^*(ds) \geq D(\xi^*). \quad (2.33)$$

На самом деле в (2.33) вместо знака неравенства стоит знак равенства. Это вытекает из того, что по определению

$$D(\xi^*) = \iint K(s, t) \xi^*(ds) \xi^*(dt),$$

и поэтому в точках t^* плана $\xi^*(dt)$ выполнено равенство

$$\int K(s, t^*) \xi^*(ds) = D(\xi^*).$$

Итак, нами доказана теорема, которую по аналогии с теоремами § 2 гл. 3 можно назвать теоремой эквивалентности.

Теорема 2.3 (теорема эквивалентности). *Необходимым и достаточным условием оптимальности плана ξ^* (т. е. того, что $\xi^* = \arg \min_{\xi \in \Xi} D(\xi)$) является выполнение равенства*

$$\inf_{t \in T} \int K(s, t) \xi^*(ds) = D(\xi^*). \quad (2.34)$$

Аналогично § 4 гл. 3 можно построить алгоритмы численного нахождения оптимальных планов. Аналогом алгоритма 4.1 из гл. 3 является следующий

Алгоритм 2.1.

1) Выбираем некоторый план ξ_0 (возможно, сосредоточенный в одной точке) и полагаем $i = 0$.

2) План ξ_{i+1} строим по формуле

$$\xi_{i+1} = (1 - \alpha_{i+1}) \xi_i + \alpha_{i+1} \xi_{(i)} \quad (2.35)$$

где $\xi_{(i)} = \xi(t_{(i)})$ — план, сосредоточенный в точке

$$t_{(i)} = \arg \min_{t \in T} \int K(s, t) \xi_i(ds). \quad (2.36)$$

3) Переходим к шагу 2 с заменой i на $i + 1$.

Основанием для выбора в (2.35) плана $\xi_{(i)} = \xi(t_{(i)})$, где $t_{(i)}$ определяется по (2.36), является то, что при таком выборе плана $\xi_{(i)}$ производная

$$\left. \frac{\partial D((1 - \alpha) \xi_i + \alpha \xi_{(i)})}{\partial \alpha} \right|_{\alpha=0+}$$

минимальна и поэтому функционал D в направлении плана $\xi_{(i)}$ убывает наиболее быстро.

В качестве последовательности $\{\alpha_{i+1}\}$ в (2.35) естественно выбирать одну из двух: либо

$$\alpha_{i+1} = 1/(i + 2), \quad (2.37)$$

либо

$$\alpha_{i+1} = \arg \min_{\alpha \in [0, 1]} D((1 - \alpha) \xi_i + \alpha \xi_{(i)}). \quad (2.38)$$

Если в алгоритме 2.1 план ξ_0 сосредоточен в одной точке, а α_{i+1} выбираются по (2.37), то планы ξ_i являются дискретными — веса входящих в них точек равны $1/(i+1)$. Способ выбора (2.38) соответствует максимально возможному уменьшению критерия D за счет выбора α_{i+1} . Величины α_{i+1} из (2.38) легко выражаются в явном виде. Действительно, из формулы, полученной в доказательстве леммы 2.1, имеем

$$D((1-\alpha)\xi_i + \alpha\xi(t_{(i)})) = (1-\alpha)^2 D(\xi_i) + \\ + 2\alpha(1-\alpha) \int K(t, t_{(i)}) \xi_i(dt) + \alpha^2 K(t_{(i)}, t_{(i)}) = a + 2b\alpha + c\alpha^2,$$

где

$$a = D(\xi_i), \quad b = \int K(t, t_{(i)}) \xi_i(dt) - D(\xi_i), \\ c = D(\xi_i) + K(t_{(i)}, t_{(i)}) - 2 \int K(t, t_{(i)}) \xi_i(dt).$$

Минимум по α квадратного трехчлена $a + 2b\alpha + c\alpha^2$ достигается при $\alpha = -b/c$ и равен $a - 3b^2/c$ (в данном случае, как нетрудно проверить (проверьте!), $b \leq 0$, $c > 0$, $|b| \leq c$). Итак, формула (2.38) эквивалентна формуле

$$\alpha_{i+1} = \frac{D(\xi_i) - \int K(t, t_{(i)}) \xi_i(dt)}{D(\xi_i) + K(t_{(i)}, t_{(i)}) - 2 \int K(t, t_{(i)}) \xi_i(dt)}. \quad (2.39)$$

2.3. Планирование эксперимента при использовании оценок типа гребневой. Пусть имеется схема регрессии $(Y, F\theta, \sigma^2 I_N)$, рассматривавшаяся в гл. 1. Как указано в гл. 1, кроме НЛН-оценок для θ , полученных с помощью метода наименьших квадратов, часто имеет смысл использовать и оценки вида

$$\hat{\theta}_N = (F^T F + A)^{-1} F^T Y. \quad (2.40)$$

Приведем примеры оценок, имеющих такой вид. Для гребневой оценки (1.59) из гл. 1 матрица A пропорциональна единичной. Если априорная информация о параметрах θ регрессионной модели $(Y, F\theta, \sigma^2 I_N)$ имеет вид

$$\theta \in \Omega = \{\theta \in \mathbb{R}^m \mid \theta^T T \theta \leq k\},$$

то минимаксная оценка (2.5) из гл. 1 выражается по формуле (2.40) с $A = \frac{\sigma^2}{k} T$. Если, наконец, на параметрическом множестве Ω задано априорное распределение $P(d\theta)$ с нулевым вектором средних и дисперсионной матрицей $V > 0$, то байесовская оценка (2.14) из гл. 1 также записывается в виде (2.40), но с $A = \sigma^2 V^{-1}$.

Отметим, что во всех указанных случаях матрица A положительно определена и известна, возможно, с точностью до скалярного множителя σ^2 .

Оценка (2.40) является смещенной.

Действительно,

$$E\hat{\theta}_N = (F^T F + A)^{-1} F^T EY = (F^T F + A)^{-1} F^T F\theta - \theta - (F^T F + A)^{-1} A\theta,$$

и, следовательно, $E\hat{\theta}_N \neq \theta$ при $\theta \neq 0$. Дисперсионная матрица оценки (2.40) равна

$$\begin{aligned} D\hat{\theta}_N &= E(\hat{\theta}_N - E\hat{\theta}_N)(\hat{\theta}_N - E\hat{\theta}_N)^T = \\ &= E[(F^T F + A)^{-1} F^T (Y - F\theta)(Y - F\theta)^T F (F^T F + A)^{-1}] = \\ &= \sigma^2 (F^T F + A)^{-1} F^T F (F^T F + A)^{-1} = \\ &= \sigma^2 (F^T F + A)^{-1} - \sigma^2 (F^T F + A)^{-1} A (F^T F + A)^{-1}. \end{aligned}$$

Рассмотрим схему регрессионного эксперимента, описанную в § 1 гл. 3. Для этой схемы оценка (2.40) строится в предположении, что план эксперимента предназначен для проведения N измерений и имеет вид

$$\xi_N = \left\{ \begin{matrix} x_1, \dots, x_N \\ 1/N, \dots, 1/N \end{matrix} \right\}. \quad (2.41)$$

где $x_1, \dots, x_N$ — некоторые точки из множества планирования X . Информационная матрица плана (2.41) равна

$$M(\xi_N) = \frac{1}{N} F^T F.$$

В терминах этой матрицы выражение для $D\hat{\theta}_N$ переписывается в виде

$$D\hat{\theta}_N = \frac{\sigma^2}{N} (M(\xi_N) + B)^{-1} - \frac{\sigma^2}{N^2} (M(\xi_N) + B)^{-1} A (M(\xi_N) + B)^{-1}, \quad (2.42)$$

где $B = \frac{1}{N} A$. При больших, но конечных N (этот случай мы и будем рассматривать ниже) матрица (2.42) почти полностью определяется матрицей

$$D_B(\xi_N) = (M(\xi_N) + B)^{-1}, \quad (2.43)$$

которую мы будем рассматривать как основу для построения критериев оптимальности планов, пренебрегая смещением оценки.

Матрица (2.43) определена для любых непрерывных планов ξ :

$$D_B(\xi) = (M(\xi) + B)^{-1}, \quad (2.44)$$

где

$$M(\xi) = \int_X f(x) f^T(x) \xi(dx) \quad (2.45)$$

есть информационная матрица плана ξ , свойства которой изучены в гл. 3. Обратная от матрицы (2.44) равна

$$M_B(\xi) = M(\xi) + B \quad (2.46)$$

и в литературе называется *байесовской информационной мат-*

рицей (по той причине, что задачи планирования чаще всего рассматривались в предположении, что оценка (2.40) является байесовской).

Ниже рассматривается задача минимизации в множестве непрерывных планов Ξ выпуклого функционала

$$\Psi_B(\xi) = \Psi[M_B(\xi)],$$

где $M_B(\xi)$ определяется по формуле (2.46), а матрица $B > 0$ фиксирована. При этом следует иметь в виду, что, поскольку матрица B зависит от числа планируемых измерений N , от этого числа будут зависеть и оптимальные непрерывные планы.

Изложенные в п. 1.4 гл. 3 свойства информационных матриц $M(\xi)$ переносятся на байесовские информационные матрицы вида (2.46) с единственным изменением: байесовские информационные матрицы $M_B(\xi)$ всегда невырождены (поскольку $B > 0$). Отсюда следует, в частности, что оптимальные по некоторым критериям планы могут быть сосредоточены в одной точке.

Критерии оптимальности планов Ψ могут быть выбраны те же, что и в гл. 3. Очевидно, что выпуклость функционала Ψ на множестве неотрицательно определенных матриц влечет его выпуклость на множестве $\mathfrak{M}_B = \{M_B(\xi), \xi \in \Xi\}$.

Предположим, что выбранный функционал $\Psi[M]$ является дифференцируемым по элементам матрицы M . Для доказательства теорем эквивалентности типа теорем из § 2 гл. 3 можно использовать описанную в указанном параграфе стандартную технику. Единственным отличием рассматриваемой в данном пункте задачи является несколько иной вид производной

$$\Delta_B(\xi_1, \xi_2) = \lim_{\alpha \rightarrow 0+} \frac{\Psi_B[(1-\alpha)\xi_1 + \alpha\xi_2] - \Psi_B[\xi_1]}{\alpha}$$

функционала Ψ_B по направлению $\xi_2 - \xi_1$ в точке $\xi_1 \in \Xi$.

Необходимое и достаточное условие оптимальности плана ξ^* , аналогичное условию (2.1) из гл. 3, имеет вид

$$\inf_{\xi \in \Xi} \Delta_B(\xi^*, \xi) \geq 0. \quad (2.47)$$

Вычисляя производную $\Delta_B(\xi^*, \xi)$, имеем

$$\begin{aligned} \Delta_B(\xi^*, \xi) &= \left. \frac{\partial \Psi[(1-\alpha)M_B(\xi^*) + \alpha M_B(\xi)]}{\partial \alpha} \right|_{\alpha=0} = \\ &= \left. \frac{\partial \Psi[(1-\alpha)M(\xi^*) + \alpha M(\xi) + B]}{\partial \alpha} \right|_{\alpha=0} = \\ &= \text{tr}[M(\xi) - M(\xi^*)] \dot{\Psi}[M_B(\xi^*)] = \\ &= \text{tr} M(\xi) \dot{\Psi}[M_B(\xi^*)] - \text{tr} M(\xi^*) \dot{\Psi}[M_B(\xi^*)] = \\ &= \int \Psi_B(x, \xi^*) \xi(dx) - \text{tr} M(\xi^*) \dot{\Psi}[M_B(\xi^*)], \end{aligned}$$

где

$$\begin{aligned}\psi_B(x, \xi^*) &= f^T(x) \hat{\Psi}[M_B(\xi^*)] f(x), \\ \hat{\Psi}[M_B(\xi^*)] &= \frac{\partial \Psi[M]}{\partial M} \Big|_{M=M_B(\xi^*)}.\end{aligned}\quad (2.48)$$

Следовательно,

$$\begin{aligned}\inf_{\xi \in \Xi} \Delta_B(\xi^*, \xi) &= \inf_{\xi \in \Xi} \int \psi_B(x, \xi^*) \xi(dx) - \\ &- \operatorname{tr} M(\xi^*) \hat{\Psi}[M_B(\xi^*)] = \inf_{x \in X} \psi_B(x, \xi^*) - \operatorname{tr} M(\xi^*) \hat{\Psi}[M_B(\xi^*)].\end{aligned}$$

Таким образом, из (2.47) вытекает справедливость следующей теоремы, аналогично теореме 2.1 из гл. 3.

Теорема 2.4 (теорема эквивалентности). Пусть Ψ — выпуклый дифференцируемый функционал на множестве неотрицательно определенных $(n \times m)$ -матриц. Тогда необходимым и достаточным условием того, что

$$\xi^* = \operatorname{argmin}_{\xi \in \Xi} \Psi[M_B(\xi)], \quad (2.49)$$

является выполнение для всех $x \in X$ неравенства

$$\psi_B(x, \xi^*) \geq \operatorname{tr} M(\xi^*) \hat{\Psi}[M_B(\xi^*)]. \quad (2.50)$$

Из этой теоремы можно получить аналоги теорем эквивалентности для различных критериев оптимальности. Для примера рассмотрим критерий D -оптимальности

$$\Psi[M] = -\ln \det M.$$

В этом случае, используя результаты п. 2.2 гл. 3, получаем

$$\begin{aligned}\psi_B(x, \xi^*) &= -f^T(x) [M_B(\xi^*)]^{-1} f(x), \\ \operatorname{tr} M(\xi^*) \hat{\Psi}[M_B(\xi^*)] &= -\operatorname{tr} M(\xi^*) [M_B(\xi^*)]^{-1};\end{aligned}$$

следовательно, необходимым и достаточным условием того, что

$$\xi^* = \operatorname{argmax}_{\xi \in \Xi} \det [M(\xi) + B],$$

является выполнение неравенства

$$f^T(x) [M(\xi^*) + B]^{-1} f(x) \leq \operatorname{tr} M(\xi^*) [M(\xi^*) + B]^{-1} \quad (2.51)$$

для всех $x \in X$.

Отметим, что в отличие от классического случая, получающегося при $B=0$, правая часть неравенства (2.51) не равна числу неизвестных параметров m .

Упражнения.

1. Используя формулу (1.9) из приложения 1, показать, что при добавлении новых результатов измерений дисперсия (2.14) оценки θ_m , рассмотренной в п. 2.2, не возрастает.
2. Докажите справедливость формулы (2.15).

3. Выведите явный вид функции (2.21) для случая (2.22) и

$$K(s, t) = \int_0^{|t-s|^{-1}} [1 - x |t - s|] e^{-x} dx.$$

4. Пусть $T = [0, 1]$, $K(s, t) = 0$, $s \neq t$. Найдите план, оптимальный по критерию (2.27).

5. Пусть $T = [0, 1]$, $K(s, t) = |t - s|$. Сравните по критерию (2.27) два плана: равномерный $\xi(dt) = dt$ и сосредоточенный с равными весами в точках нуля и единица. Проверьте оптимальность этих планов с помощью теоремы 2.2.

6. Покажите, что величины α_{i+1} из (2.39) удовлетворяют неравенству $0 \leq \alpha_{i+1} \leq 1$.

7. Сформулируйте модификацию алгоритма 2.1, в которой аналогично алгоритму 4.2 из гл. 3 допускаются отрицательные значения величин α_{i+1} . Выведите формулу для отрицательных значений α_{i+1} , аналогичную формуле (2.39).

8. Покажите, что оценка (2.40) является асимптотически несмещенной.

9. Покажите, что матрица вторых моментов отклонений оценки (2.40) от истинных значений параметров θ равна

$$E(\hat{\theta}_N - \theta)(\hat{\theta}_N - \theta)^T = (F^T F + A)^{-1} + (F^T F + A)^{-1} A (\theta \theta^T A - I_m) (F^T F + A)^{-1}.$$

Выведите отсюда, что эта матрица при больших значениях N близка к матрице $(F^T F + A)^{-1}$.

10. Докажите, что для любого непрерывного плана ξ выполнено равенство

$$\int_X \psi_B(x, \xi) \xi(dx) = \text{tr } M(\xi) \Psi^{\alpha}[M_B(\xi)].$$

11. Используя результат предыдущего упражнения, покажите, что в точках плана (2.49) неравенство (2.50) превращается в равенство.

12. Покажите, что для критерия L -оптимальности $\Psi[M] = \text{tr } LM^{-1}$ функция $\psi_B(x, \xi)$ имеет вид

$$\psi_B(x, \xi) = f^T(x) [M_B(\xi)]^{-1} L [M_B(\xi)]^{-1} f(x),$$

а правая часть неравенства (2.50) — вид

$$\text{tr } L [M_B(\xi^*)]^{-1} M(\xi^*) [M_B(\xi^*)]^{-1}.$$

13. Используя результаты упражнений 11 и 12, сформулируйте теорему эквивалентности для байесовских L -оптимальных планов.

14. Сформулируйте численные методы построения оптимальных байесовских планов (2.49), аналогичные алгоритмам 4.1—4.3 из гл. 3.

15. Пусть функция регрессии имеет вид $\eta(x) = \theta_1 + \theta_2 x$, $X = [-1, 1]$, а матрица B пропорциональна единичной: $B = cI_2$, $c > 0$. По критериям байесовской D -оптимальности и байесовской A -оптимальности сравните планы

$$\xi_1(dx) = \frac{1}{2} dx, \quad \xi_2 = \left\{ \begin{matrix} -1 & 1 \\ 1/2 & 1/2 \end{matrix} \right\}, \quad \xi_3 = \left\{ \begin{matrix} -1 & 0 & 1 \\ 1/3 & 1/3 & 1/3 \end{matrix} \right\}, \quad \xi_4 = \left\{ \begin{matrix} 0 \\ 1 \end{matrix} \right\}.$$

С помощью теоремы эквивалентности проверьте указанные планы на оптимальность.

3.3. Планирование эксперимента при восстановлении линейного оператора. Предположим, что истинный оператор L_n , определяемый изучаемым процессом, неизвестен, но известен близкий к нему оператор L , для которого может быть точно решена основная задача (3.1), а для любых $p_i \in L^2(X, v)$ — сопряженные задачи

$$L^* \varphi_{p_i}^* = p_i, \quad i = 1, 2, \dots$$

Предположим также, что для любых $p_i \in L^2(X, v)$ ($i = 1, 2, \dots$) могут быть со случайной ошибкой измерены функционалы $\mathcal{J}_p[\varphi_n]$ от решения основного уравнения

$$L_n \varphi_n = q;$$

результат измерения для функции p_i — случайная величина

$$v_i = (\varphi_n, p_i) + \varepsilon_i = \mathcal{J}_{p_i}[\varphi_n] + \varepsilon_i,$$

где случайные ошибки $\varepsilon_1, \varepsilon_2, \dots$ взаимно независимы, одинаково распределены, имеют нулевое среднее и конечную дисперсию σ^2 .

Положим $L' = L_n$, $\varphi' = \varphi_n$. Приращение оператора δL будем искать в параметрическом виде: предположим, что

$$\delta L = \sum_{i=1}^m \theta_i A_i, \quad (3.9)$$

где A_i ($i = 1, \dots, m$) — известные линейные операторы, $\theta_1, \dots, \theta_m$ — неизвестные параметры, подлежащие оцениванию.

Задача планирования эксперимента заключается в таком выборе функций $\tilde{p}_1, \dots, \tilde{p}_N$ (N — число измерений), чтобы неизвестные параметры $\theta = (\theta_1, \dots, \theta_m)'$ были оценены наиболее точно.

По формуле малых возмущений (3.8) имеем

$$(\varphi_{p_i}^*, \delta L \varphi) = -\delta \mathcal{J}_{p_i}, \quad i = 1, \dots, N. \quad (3.10)$$

Используя (3.9) и то, что функционалы $\mathcal{J}_{p_i}[\varphi_n]$ вычисляются со случайной ошибкой, из (3.10) имеем

$$\sum_{i=1}^m \theta_i (\varphi_{p_i}^*, A_i \varphi) = -\delta \mathcal{J}_{p_i} = -E\{v_i - \mathcal{J}_{p_i}\}. \quad (3.11)$$

Обозначив $f_i(x) = A_i \varphi(x)$ ($i = 1, \dots, m$), замечаем, что рассматриваемая задача планирования эксперимента по оцениванию параметров θ_i может быть рассмотрена с общей позиции планирования регрессионного эксперимента в функциональном пространстве (см. п. 2.1). Приведем простейший путь решения этой задачи.

Положим

$$F = \left[(\varphi_{p_i}^*, A_j \varphi) \right]_{\substack{i=1, \dots, N \\ j=1, \dots, m}}, \quad Y = \begin{bmatrix} \mathcal{J}_{p_1} - v_1 \\ \vdots \\ \mathcal{J}_{p_N} - v_N \end{bmatrix} \quad (3.12)$$

и перепишем уравнение (3.11) в матричном виде $EY = F\theta$. Учитывая условия на ошибки измерений, пишем $Y \sim (F\theta, \sigma^2 I_N)$ и, следовательно, приходим к классической линейной регрессионной модели $(Y, F\theta, \sigma^2 I_N)$ (см. п. 1.2 гл. 1).

Выбирая в качестве метода оценивания метод наименьших квадратов, получаем, согласно (1.20) из гл. 1, оценку

$$\hat{\theta} = (F^T F)^{-1} F^T Y$$

для неизвестных параметров θ . Информационная матрица $F^T F$ в рассматриваемом случае равна

$$F^T F = \left| \sum_{i=1}^N (\varphi_{p_i}^*, A_i \varphi) (\varphi_{p_i}^*, A_i \varphi) \right|_{L, l=1}^m.$$

Следовательно, нормированная информационная матрица $M(\xi)$ произвольного плана (т. е. произвольной вероятностной меры на $L^2(X, \nu)$) равна

$$M(\xi) = \left| \int (\varphi_p^*, A_l \varphi) (\varphi_p^*, A_l \varphi) \xi(d\varphi_p^*) \right|_{L, l=1}^m.$$

Так же, как и в п. 2.1, посредством преобразования

$$z_l = (\varphi_p^*, A_l \varphi)$$

задача планирования приводится к задаче планирования для линейной модели

$$Ey(z) = \sum_{i=1}^m \theta_i z_i, \quad z \in Z \subset \mathbb{R}^m.$$

Формулируя критерий оптимальности плана и находя с помощью результатов гл. 3, 4 оптимальный план, на следующем этапе с помощью процедуры п. 2.1 восстанавливаем функции $\tilde{\varphi}_{p_1}^*, \dots, \tilde{\varphi}_{p_N}^*$, соответствующие N точкам дискретного оптимального плана.

Наконец, замечаем, что нашей целью являлся оптимальный выбор не функций $\tilde{\varphi}_{p_i}^*$, а функций $\tilde{p}_1, \dots, \tilde{p}_N$; оптимальные функции $\tilde{p}_i$ легко находим через оптимальные функции $\tilde{\varphi}_{p_i}^*$ по формуле

$$\tilde{p}_i = L^* \tilde{\varphi}_{p_i}^*.$$

Выше рассмотрен случай, когда решение модельной задачи близко к решению реальной, т. е. когда можно пользоваться формулой малых возмущений (3.8). Если точность формулы (3.8) мала, то получаемую после проведения указанных выше вычислений оценку истинного оператора можно считать лишь первым приближением к решению задачи его восстановления, используя эту оценку при аналогичных вычислениях в качестве невозмущенного оператора. Проведение такого рода последовательных вычислений соответствует применению последовательного подхода (см. § 3 гл. 5).

Методы планирования эксперимента, основанные на теории возмущений операторов, применимы также и в том случае, когда априорная информация о функции $\varphi(x)$ задается уравнением

$$A\varphi = q, \quad (3.13)$$

где A — нелинейный оператор. Рассмотрим этот случай, который не менее важен в практическом отношении.

Пусть F_1 (множество, на котором определен оператор A) и F_2 (множество значений A) есть подмножества линейного нормированного пространства и оператор A имеет производную Фреше в окрестности интересующего нас решения φ уравнения (3.13). Из существования производной Фреше A'_φ оператора A следует, в частности, что малым по норме возмущениям правой части q соответствуют малые по норме возмущения φ .

По определению производной Фреше справедливо приближенное равенство $A(\varphi + \delta\varphi) - A\varphi \approx A'_\varphi \delta\varphi$ (т. е. $A'_\varphi \delta\varphi \approx \delta q$), которое обычно рассматривают как точное, пренебрегая величинами более высокого по норме порядка малости, чем $\delta\varphi$. Поскольку A'_φ — линейный при фиксированном φ оператор, то мы приходим к линейному случаю.

Дополнительным требованием для получения содержательных результатов является непрерывная (по норме) зависимость A'_φ от φ , что, безусловно, имеет место, если существует вторая производная Фреше A''_φ . Выполнение условия непрерывности позволяет в приближенном равенстве

$$A'_\varphi \delta\varphi \approx \delta q$$

использовать значение A'_φ для приближенного значения $\hat{\varphi}$, а не для точного φ . Это очень важно в практических задачах.

Далее можно ввести в рассмотрение, как и в линейном случае, сопряженное уравнение

$$(A'_\varphi)^* \varphi^*[\varphi, h] = h. \quad (3.14)$$

Здесь обозначение $\varphi^*[\varphi, h]$ для решения сопряженного уравнения подчеркивает его зависимость как от φ , так и от h . Дальнейшее применение идей последовательного подхода не требует специальных пояснений.

Если $A = A(\varphi, \alpha)$, $q = q(\beta)$ (где $\alpha = (\alpha_1, \dots, \alpha_n)^T$, $\beta = (\beta_1, \dots, \beta_m)^T$ — параметры из некоторых параметрических множеств) и соответствующие производные от A и q существуют и непрерывно зависят от этих параметров, то равенство $A\varphi = q$ для возмущенной системы часто заменяют приближенным равенством

$$\bar{A}\varphi + \sum_{i=1}^n \left(\frac{\partial \bar{A}}{\partial \alpha_i} \right) \delta \alpha_i \varphi = \bar{q} + \sum_{j=1}^m \frac{\partial \bar{q}}{\partial \beta_j} \delta \beta_j, \quad (3.15)$$

где $\alpha_i = \bar{\alpha}_i + \delta\alpha_i$, $\beta_i = \bar{\beta}_i + \delta\beta_i$, а черта над символом соответствует невозмущенному состоянию системы. Так как $\bar{A}\bar{\varphi} = \bar{q}$, то приближенное равенство (3.15) может быть очевидным образом упрощено.

Упражнения.

1. Пусть $X = [0, 1]$, а

$$A = -\frac{d}{dx} a(x) \frac{d}{dx} + b(x)$$

есть оператор, определенный для таких дважды дифференцируемых функций $\varphi(x)$, что $\varphi(0) = \varphi(1) = 0$:

$$a(x) = \bar{a}(x) + a_1 u_1(x) + a_2 u_2(x),$$

$$b(x) = \bar{b}(x) + b_1 v_1(x) + b_2 v_2(x)$$

где $\bar{a}$, u_1 , u_2 , $\bar{b}$, v_1 , v_2 — известные функции, a_1 , a_2 , b_1 , b_2 — неизвестные параметры. Эксперимент состоит в измерении функционалов

$$(p_i, \bar{\varphi}) = \int_X p_i(x) \varphi(x) dx$$

где p_i ($i = 1, \dots, N$) заданы. $\bar{\varphi}$ — решение уравнения $\bar{A}\bar{\varphi} = q$, $\bar{\varphi}(0) = \bar{\varphi}(1) = 0$, функция q задана. $\bar{A}$ — невозмущенный оператор (оператор A при $a_1 = a_2 = b_1 = b_2 = 0$).

Запишите для рассматриваемого случая уравнения вида (3.4) при $a(x) = 1$, $u_1 = x$, $u_2 = x^2$, $\bar{b}(x) = 0$, $v_1 = 1$, $v_2 = x$, $p_i = x^i$ ($i = 1, 2, \dots, 6$), $q = 1$.

Вычислите при тех же данных элементы матрицы F (см. (3.12)).

2. Пусть $A\varphi = -d^2\varphi/dx^2 + (a(x)\varphi + b(x))\varphi$, $\varphi(0) = \varphi(1) = 0$, $a(x) = 1 + a_1 x$, $b(x) = 1 + b_1 x$, где a_1 и b_1 — неизвестные параметры. Измеряются скалярные произведения $(x^i, \bar{\varphi})$. $\bar{\varphi}$ — решение уравнения $-d^2\bar{\varphi}/dx^2 + \bar{\varphi}^2 + \varphi = 1$.

Напишите для рассматриваемого случая уравнение вида (3.13).

Напишите выражение для матрицы F при линеаризации оператора A .

§ 4. Планирование эксперимента при неадекватности линейной модели

4.1. Постановка задач. В классической постановке задачи планирования регрессионного эксперимента предполагается, что функция регрессии $\eta(x) = Ey(x)$ является линейной комбинацией базисных функций $f_1(x), \dots, f_m(x)$, т. е. представима в виде

$$\eta(x) = \theta^T f(x), \quad (4.1)$$

где $\theta = (\theta_1, \dots, \theta_m)^T$ — вектор неизвестных параметров, $f(x) = (f_1(x), \dots, f_m(x))^T$. На практике более реальной является ситуация, когда представление (4.1) выполняется не точно, а с некоторой погрешностью, т. е. вместо (4.1) имеет место равенство

$$\eta(x) = \theta^T f(x) + \psi(x), \quad (4.2)$$

где неизвестная функция $\psi(x)$ априори принадлежит некоторому множеству Ψ , причем обычно предполагают, что множество Ψ состоит из функций, в каком-то смысле не слишком отличающихся от нуля. Задачи планирования для случая (4.2) принято называть *устойчивыми (робастными)* по отношению к предположению о справедливости модели (4.1). Планы, оптимальные для наихудшей функции ψ из множества Ψ , называют *робастными планами*.

Следующий пример показывает, что применение в ситуации (4.2) планов, оптимальных для ситуации (4.1), может приводить к неудовлетворительным результатам.

Пример 4.1. Пусть множество планирования X — отрезок $[-1, 1]$, $m = 2$, $Dy(x) = \sigma^2$, $f_1(x) = 1$, $f_2(x) = x$,

$$\Psi \subset \{\psi(x) = ax^2, x \in [-1, 1], a \in \mathbb{R}\}.$$

Таким образом, в рассматриваемом случае модель (4.1) линейная ($\eta(x) = \theta_1 + \theta_2 x$), а модель (4.2) квадратичная ($\eta(x) = \theta_1 + \theta_2 x + \theta_3 x^2$). Согласно результатам § 3 гл. 3, для модели (4.1) оптимальным в различных смыслах является план, сосредоточенный с равными весами в точках $x_1 = 1$, $x_2 = -1$. Однако для квадратичной модели указанный план является вырожденным и не позволяет получить несмещенные оценки всех трех параметров квадратичной функции регрессии.

В данном параграфе будут рассматриваться различные конкретизации следующей общей постановки задач оценивания регрессии.

Предположим, что в точках $x_i \in X \subset \mathbb{R}^k$ ($j = 1, \dots, N$), определяемых планом проведения эксперимента, могут быть вычислены значения

$$\begin{aligned} y(x_j) &= \eta(x_j) + e(x_j), \quad E e(x_j) = 0, \\ E e^2(x) &= \sigma^2(x), \quad E e(x_i) e(x_j) = 0 \quad (i \neq j), \end{aligned} \quad (4.3)$$

причем функция регрессии $\eta(x)$ априори принадлежит некоторому множеству функций $\mathcal{F}$. Например, в случае справедливости представления (4.2) $\mathcal{F} = \{\eta | \eta(x) = \theta^T f(x) + \psi(x), \psi \in \Psi\}$. Если множество функций $\mathcal{F}$ бесконечномерно, то задача оценивания функции регрессии $\eta(x)$ называется *задачей непараметрического оценивания* регрессии. В качестве оценок $\hat{\eta}(x)$ функции регрессии $\eta(x)$ будем рассматривать только линейные оценки, т. е. оценки, представимые в виде

$$\hat{\eta}(x) = \hat{\eta}_N(x) = \sum_{j=1}^N c_{j,N}(x) y(x_j), \quad (4.4)$$

где коэффициенты $c_{j,N}(x)$ ($j = 1, \dots, N$) не зависят от результатов измерений и неизвестной функции регрессии $\eta(x)$, но могут зависеть от $N, x, x_1, \dots, x_N$.

В качестве меры точности оценки функции регрессии $\hat{\eta}(x)$ будем использовать среднеквадратичную погрешность, квадрат

которой определяется по формуле

$$J = E \int_X (\hat{\eta}(x) - \eta(x))^2 w(x) dx. \quad (4.5)$$

Здесь $w(x)$ — неотрицательная весовая функция, которую вводят при необходимости придать различную значимость ошибкам приближения функции регрессии в различных точках множества X . Ничего не изменится, если в (4.5) вместо меры $w(x)dx$ использовать произвольную σ -конечную меру $\nu(dx)$ на σ -алгебре $\mathcal{B}$, т. е. записывать J в виде

$$J = E \int_X (\hat{\eta}(x) - \eta(x))^2 \nu(dx).$$

Преобразуем выражение (4.5).

Лемма 4.1. *Предположим, что конечны величины*

$$B = \int_X (\eta(x) - E\hat{\eta}(x))^2 w(x) dx, \quad (4.6)$$

$$V = E \int_X (E\hat{\eta}(x) - \hat{\eta}(x))^2 w(x) dx. \quad (4.7)$$

Тогда величина J , определенная по (4.5), представима в виде суммы

$$J = B + V. \quad (4.8)$$

Доказательство. Имеем следующее разложение погрешности $\hat{\eta}(x) - \eta(x)$:

$$\eta(x) - \hat{\eta}(x) = [\eta(x) - E\hat{\eta}(x)] + [E\hat{\eta}(x) - \hat{\eta}(x)].$$

Возводя обе части этого равенства в квадрат, получаем

$$\begin{aligned} [\eta(x) - \hat{\eta}(x)]^2 &= [\eta(x) - E\hat{\eta}(x)]^2 + \\ &+ 2[\eta(x) - E\hat{\eta}(x)][E\hat{\eta}(x) - \hat{\eta}(x)] + [E\hat{\eta}(x) - \eta(x)]^2. \end{aligned}$$

Полученное равенство интегрируем по мере $w(x)dx$ и берем математическое ожидание. Соотношение (4.8) вытекает из того, что

$$\begin{aligned} E \int_X [\eta(x) - E\hat{\eta}(x)][E\hat{\eta}(x) - \hat{\eta}(x)] w(x) dx &= \\ &= \int_X E \{[\eta(x) - E\hat{\eta}(x)][E\hat{\eta}(x) - \hat{\eta}(x)]\} w(x) dx = \\ &= \int_X [\eta(x) - E\hat{\eta}(x)][E\hat{\eta}(x) - E\hat{\eta}(x)] w(x) dx = 0. \end{aligned}$$

Согласно теореме Фубини, перестановка интеграла и математического ожидания в приведенной цепочке равенств возможна, если конечна величина

$$\tilde{J} = E \int_X |E[\eta(x) - E\hat{\eta}(x)][E\hat{\eta}(x) - \hat{\eta}(x)]| w(x) dx.$$

В силу неравенства Коши — Буняковского имеем $J^2 \leq BV$, и поэтому величина J конечна. Лемма доказана.

Величина $\sqrt{B}$, где B определяется по формуле (4.6), называется *систематической погрешностью* и показывает, насколько в метрике пространства L_2 математическое ожидание оценки $\hat{\eta}(x)$ удалено от истинной функции регрессии $\eta(x)$. Обозначение квадрата систематической погрешности символом B объясняется тем, что указанная величина представляет собой меру смещения оценки (от английского эквивалента термина «смещение» — bias).

Величина $\sqrt{V}$, где V определяется по формуле (4.7), называется *случайной погрешностью* и является характеристикой, аналогичной обычной дисперсии (по-английски «дисперсия» — это variance, отсюда и символ V). Более того, поскольку при фиксированном $x \in X$ дисперсия оценки $\hat{\eta}(x)$ равна

$$D\hat{\eta}(x) = E(E\hat{\eta}(x) - \hat{\eta}(x))^2,$$

то квадрат случайной погрешности V может быть записан в виде

$$V = \int_X [D\hat{\eta}(x)] w(x) dx. \quad (4.9)$$

Величина $\sqrt{J}$, где J определяется по формуле (4.5), называется *суммарной погрешностью*.

В рассмотренной ситуации возможны различные постановки задачи планирования эксперимента. Эти постановки зависят от вида рассматриваемой модели регрессии (т. е. от множества функций $\mathcal{F}$), от способа оценивания функции регрессии и от критерия оптимальности, который обычно конструируется путем смешивания погрешностей B и V или предпочтения одной из них.

Постановки задач оптимального планирования при неадекватности линейной модели, рассмотренные в этом параграфе, не охватывают всего их разнообразия, но дают о них некоторое представление. Существенно иной класс задач планирования получается в случае, когда функция регрессии известна с точностью до нелинейно входящих в нее параметров (этому классу задач посвящена гл. 5).

4.2. Оптимальное планирование в конечномерных пространствах функций. Рассмотрим случай, когда $\sigma^2(x) = \sigma^2$, функция регрессии $\eta(x)$ представима в виде (4.2) и множество функций $\Psi = \{\psi\}$ конечномерно. Пусть множество Ψ порождается функциями $g_1(x), \dots, g_l(x)$, т. е. для любой $\psi \in \Psi$ можно найти такие числа $\beta_1, \dots, \beta_l$, что

$$\psi(x) = \sum_{i=1}^l \beta_i g_i(x) = \beta^T g(x), \quad (4.10)$$

где

$$\beta = (\beta_1, \dots, \beta_l)^T, \quad g(x) = (g_1(x), \dots, g_l(x))^T,$$

причем будем считать, что функции

$$f_1, \dots, f_m, g_1, \dots, g_l$$

ортонормированы с весом $\omega(x)$. Согласно (4.2) и (4.10), для истинной функции регрессии $\eta(x)$ справедливо представление

$$\eta(x) = \theta^T f(x) + \beta^T g(x), \quad (4.11)$$

где θ и β — векторы неизвестных параметров. Предположим, что в силу какой-либо причины исследователь вынужден вместо оценки функции регрессии (4.11) ограничиться оценкой функции $\eta_1(x) = \theta^T f(x)$ (например, в силу невозможности оценивания всех параметров функции (4.11)). В качестве оценки $\hat{\eta}_1(x)$ для $\eta_1(x)$ выбирается, таким образом,

$$\hat{\eta}_1(x) = \hat{\theta}^T f(x), \quad (4.12)$$

где $\hat{\theta}$ — некоторая оценка параметров θ , построенная по результатам измерений вида (4.3).

Пусть $x_1, \dots, x_N$ — точки проведения измерений (не обязательно различные), $Y = (y(x_1), \dots, y(x_N))^T$ — вектор результатов измерений. В силу линейности метода оценивания статистика $\hat{\theta}$ записывается в виде

$$\hat{\theta} = AY, \quad (4.13)$$

где A — некоторая $m \times N$ -матрица, определяемая способом оценивания. По лемме 1.3 из гл. 1 дисперсионная матрица оценок $\hat{\theta}$ равна

$$D\hat{\theta} = ADYA^T = \sigma^2 AA^T. \quad (4.14)$$

Запишем план эксперимента в виде

$$\xi = \left\{ \begin{matrix} x_1, \dots, x_n \\ p_1, \dots, p_n \end{matrix} \right\}, \quad (4.15)$$

где $p_i = r_i/N$ ($i = 1, \dots, n$), r_i — число наблюдений в точке x_i . Тогда (4.13) и (4.14) примут вид

$$\hat{\theta} = C\bar{Y}, \quad (4.16)$$

$$D\hat{\theta} = \frac{\sigma^2}{N} CP^{-1}C^T, \quad (4.17)$$

где C — некоторая $m \times n$ -матрица (зависящая от точек $x_1, \dots, x_n$, но не зависящая от $p_1, \dots, p_n$), n -вектор $\bar{Y}$ составлен из среднеарифметических результатов измерений $y(x_i)$ в точках x_i ($i = 1, \dots, n$), а матрица P имеет вид

$$P = \begin{bmatrix} p_1 & 0 & \dots & 0 \\ 0 & p_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & p_n \end{bmatrix}.$$

Аналогично классической теории регрессионного планирования в качестве дисперсионной матрицы непрерывного плана

$$\xi = \left\{ \begin{matrix} x_1, & \dots, & x_n \\ p_1, & \dots, & p_n \end{matrix} \right\}, \quad p_i \geq 0, \quad \sum_{i=1}^n p_i = 1, \quad (4.18)$$

можно выбрать матрицу

$$D(\xi) = CP^{-1}C^T. \quad (4.19)$$

В силу (4.9), (4.17) и ортонормированности функций $f_1, \dots, f_m$ квадрат случайной погрешности V при проведении измерений согласно плану (4.15) равен

$$\begin{aligned} V &= \int_{\bar{x}} (D\hat{\eta}(x)) w(x) dx = \int_{\bar{x}} D[\hat{\theta}^T f(x)] w(x) dx = \\ &= \text{tr } D\hat{\theta} \int_{\bar{x}} f(x) f^T(x) w(x) dx = \frac{\sigma^2}{N} \text{tr } D(\xi). \end{aligned}$$

Для рассматриваемого случая квадрат систематической погрешности B равен

$$\begin{aligned} B &= \int_{\bar{x}} (\eta(x) - E\hat{\eta}(x))^2 w(x) dx = \\ &= \int_{\bar{x}} [(\theta - E\hat{\theta})^T f(x) + \beta^T g(x)]^2 w(x) dx = \\ &= \int_{\bar{x}} [(\theta - E\hat{\theta})^T f(x)]^2 w(x) dx + \int_{\bar{x}} [\beta^T g(x)]^2 w(x) dx. \end{aligned}$$

Отсюда следует, что квадрат систематической погрешности B минимален, если оценки $\hat{\theta}$ несмещенные, т. е. $E\hat{\theta} = \theta$ для всех θ . С учетом (4.16) это условие переписывается в виде

$$\theta = CE\bar{\theta} = CF\bar{\theta}, \quad (4.20)$$

где $\bar{\theta} = (\theta, \beta)^T$ — вектор размерности $m + l$,

$$F = \|f_1(x_1), \dots, f_m(x_1), g_1(x_1), \dots, g_l(x_1)\|_{l=1}^n.$$

Поскольку $\theta = T\bar{\theta}$, где $T = \|I_m, 0\|$, то выполнение равенства (4.20) для всех θ эквивалентно матричному равенству

$$CF = T, \quad (4.21)$$

которое представляет собой условие только на точки плана (4.15).

Полученные результаты позволяют сформулировать следующую задачу непрерывного оптимального планирования эксперимента, которую в литературе иногда называют *задачей несмещенного планирования*.

Предположим, что измерения (4.3) проводятся согласно непрерывному плану эксперимента (4.18), истинная функция регрессии имеет вид (4.11) и оценкой для нее является (4.12), где $\hat{\theta}$ определяется по (4.16). Выбирая из условия (4.21) точки $x_1, \dots, x_n$ плана (4.18) и метод оценивания, определяемый матрицей C , мы тем самым минимизируем систематическую погрешность $\sqrt{B}$. Оставшуюся свободу в выборе весов $p_1, \dots, p_n$ плана (4.18) используем для минимизации квадрата случайной погрешности

$$V = c \operatorname{tr} D(\xi), \quad (4.22)$$

где дисперсионная матрица плана ξ определяется по (4.19). Отметим, что рассматриваемая задача наиболее интересна, если на число точек n плана ξ наложено ограничение $n < k + l$, так как в противном случае в качестве ξ можно выбрать план, невырожденный для модели (4.11).

Покажем, что задача минимизации случайной погрешности по весам плана легко решается. Преобразуем (4.22):

$$V = c \operatorname{tr} D(\xi) = c \operatorname{tr} C P^{-1} C^T = c \operatorname{tr} (C^T C) P^{-1}.$$

Обозначим диагональные элементы матрицы $C^T C$ через b_i^2 ($i = 1, \dots, n$). Тогда

$$\operatorname{tr} (C^T C) P^{-1} = \sum_{i=1}^n \frac{b_i^2}{p_i}. \quad (4.23)$$

Из неравенства Коши (следствие 2.2 из приложения 1) следует, что

$$\sum_{i=1}^n |b_i| = \sum_{i=1}^n \frac{|b_i|}{\sqrt{p_i}} \sqrt{p_i} \leq \left(\sum_{i=1}^n \frac{b_i^2}{p_i} \right)^{1/2} \left(\sum_{i=1}^n p_i \right)^{1/2},$$

причем равенство имеет место только при

$$b_i^2/p_i = k p_i, \quad k > 0, \quad i = 1, \dots, n.$$

Отсюда вытекает, что минимальное значение выражения (4.23) на множестве

$$\left\{ p_1, \dots, p_n \mid p_i \geq 0, \sum_{i=1}^n p_i = 1 \right\}$$

достигается при

$$p_i^* = |b_i| / \sum_{i=1}^n |b_i|.$$

В принципе можно поставить задачу минимизации произвольного выпуклого функционала $\Phi[D(\xi)]$ по весам плана ξ . Эта задача является частным случаем задачи (2.4) из гл. 8.

4.3. Влияние способа смешивания погрешностей на вид оптимального плана. Из (4.6) вытекает, что квадрат систематической погрешности $B = B(\eta)$ зависит от неизвестной функции η . Поэтому вместо него обычно рассматривают величину

$$B_{\max} = \max_{\eta \in \mathcal{F}} B(\eta),$$

т. е. наихудшее значение квадрата систематической погрешности на множестве функций $\mathcal{F}$. При замене B на $B_{\max}$ квадрат суммарной погрешности J аналогично заменяется на

$$J_{\max} = B_{\max} + V. \quad (4.24)$$

Если на множестве $\mathcal{F}$ можно определить меру λ , отражающую дополнительные сведения о неизвестной функции $\eta \in \mathcal{F}$, то вместо $B_{\max}$ рассматривают величину

$$\int_{\mathcal{F}} B(\eta) \lambda(d\eta)$$

(байесовский критерий).

В отличие от систематической погрешности при любом линейном способе оценивания (4.4) случайная погрешность не зависит от неизвестной функции регрессии. Действительно, используя условия (4.3), для оценки (4.4) получаем

$$\begin{aligned} V &= \int_X E[\hat{\eta}(x) - E\hat{\eta}(x)]^2 w(x) dx = \\ &= \int_X E\left[\sum_{j=1}^N c_{j,N}(x)(y(x_j) - \eta(x_j))\right]^2 w(x) dx = \\ &= \int_X \left[\sum_{j=1}^N \sigma^2(x_j) c_{j,N}^2(x)\right] w(x) dx. \end{aligned}$$

Случайная погрешность $\sqrt{V}$ не зависит от неизвестной функции η , но зависит от дисперсий измерений, которые обычно неизвестны. Таким образом, критерий оптимального выбора линейной оценки (4.4) и плана эксперимента, определяющего точки проведения измерений, однозначно определен быть не может. Величины же $B_{\max}$ и V по отдельности с точностью до постоянного множителя обычно определяются, и поэтому могут быть выбраны в качестве критериев оптимальности.

В большинстве постановок задач планирования эксперимента при неадекватности линейной модели критерий $B_{\max}$ предпочитается критерию V . Тем самым способ оценивания и план эксперимента выбираются так, чтобы величина $B_{\max}$ была наименьшей. Условие минимизации $B_{\max}$ при этом не полностью определяет план эксперимента, и оставшаяся свобода в выборе плана используется для минимизации случайной погрешности. Получающиеся при таком подходе планы называются *робаст-*

ными или несмещенными. Происхождение последнего термина поясняется в п. 4.4.

В некоторых случаях (один из них рассмотрен в п. 4.5) систематическая погрешность полностью определяется выбранным методом оценивания, и поэтому в качестве критерия оптимальности плана эксперимента выбирается величина случайной погрешности. Существенно иной подход к конструированию критерия оптимальности содержится в п. 4.6 и состоит в требовании равенства скоростей убывания систематической и случайной погрешностей при числе измерений N , стремящемся к бесконечности.

Ниже на простом примере показано, что способ смешивания случайной и систематической погрешностей оказывает существенное влияние на вид точного оптимального плана.

Пусть $X = [0, 1]$, $w(x) = 1$, $\sigma^2(x) = \sigma^2$, оценка $\hat{\eta}(x)$ — многочлен первой степени, методом оценивания является метод наименьших квадратов, $\mathcal{F}$ — множество функций η , имеющих на $[0, 1]$ такую интегрируемую с квадратом первую производную, что

$$\max_{x \in X} |\eta'(x)| \leq M.$$

При указанных предположениях имеем для $\eta(x)$ представление

$$\eta(x) = \eta(0) + \int_0^x \eta'(t) dt,$$

или

$$\eta(x) = \eta(0) + \int_0^1 \eta'(t) E(x-t) dt,$$

где

$$E(z) = \begin{cases} 1, & \text{если } z > 0, \\ 0, & \text{если } z \leq 0. \end{cases}$$

Отсюда получаем

$$\eta(x) - \sum_{j=1}^N c_{j,N}(x) \eta(x_j) = \int_0^1 \eta'(t) K(x, t) dt,$$

где

$$K(x, t) = E(x-t) - \sum_{j=1}^N c_{j,N}(x) E(x_j - t).$$

Следовательно,

$$J \leq M^2 \int_0^1 \int_0^1 K^2(x, t) dx dt + \sigma^2 \sum_{j=1}^N \int_0^1 c_{j,N}^2(x) dx, \quad (4.25)$$

причем в множестве функций $\mathcal{F}$ существует такая (наихудшая) функция, для которой в (4.25) достигается знак равенства.

Отсюда вытекает, что в рассматриваемом случае квадрат суммарной погрешности (4.24) имеет вид

$$J_{\text{max}} = M^2 \int_0^1 \int_0^1 K^2(x, t) dx dt + \sigma^2 \sum_{i=1}^N \int_0^1 c_{i, N}^2(x) dx.$$

Если известно отношение σ^2/M^2 , то оптимальный план можно найти, минимизируя величину J_{max} по набору точек $x_1, \dots, x_N$. Если отношение σ^2/M^2 неизвестно, то задача планирования оказывается двухкритериальной.

В случае $M=0$ функция регрессии имеет вид $\eta(x) = \theta_1 + \theta_2 x$, поэтому точки оптимального плана расположены на концах промежутка $[0, 1]$ с равными весами. Для σ^2/M^2 , отличных от бесконечности, точки плана распределены по всему отрезку. Оптимальные планы для случая $N=5$ при разных σ^2/M^2 приведены в табл. 4.1.

Таблица 4.1

x_j	σ^2/M^2				
	0	0,2	0,4	0,6	0,8
x_1	0,085	0,043	0,017	0	0
x_2	0,283	0,222	0,185	0,157	0,130
x_3	0,5	0,5	0,5	0,5	0,5
x_4	0,717	0,778	0,815	0,843	0,870
x_5	0,915	0,957	0,983	1	1

4.4. Рандомизованные несмещенные планы. Развивая точку зрения непрерывного планирования, которая рассматривает план эксперимента как вероятностную меру, можно наметить подход, который получил название *несмещенного планирования*.

Пусть $\hat{\eta}(x) = \hat{\eta}(x, x_1, \dots, x_N)$ — оценка функции регрессии, построенная по результатам $y(x_1), \dots, y(x_N)$ измерений функции $y(x)$ в точках $x_1, \dots, x_N$. Будем считать $x_1, \dots, x_N$ случайным вектором, т. е. определим вероятностную меру на множестве таких векторов. Пусть эта мера определяется плотностью совместного распределения $p(x_1, \dots, x_N)$. Пусть плотность распределения p , которую будем называть планом эксперимента, выбирается так, чтобы в некоторой метрике ρ расстояние

$$\rho_p = \rho(\hat{\eta}(x), E\hat{\eta}(x, x_1, \dots, x_N))$$

было минимальным (при этом символ математического ожидания означает осреднение по плану эксперимента как по вероятностной мере, так и по выборочному пространству, в рамках которого производится эксперимент). План эксперимента p_* , для которого расстояние ρ_{p_*} минимально, будем называть *несмещенным планом* (несмещенным по отношению к элементу наилучшего приближения в метрике ρ).

Сравнительно простые результаты удается получить для случая, когда метрика ρ является квадратичной (в этом случае $\rho_p = \sqrt{B_{\max}}$).

Пусть $L^2(X, \mu)$ — пространство функций, интегрируемых с квадратом на множестве X относительно вероятностной меры μ , и пусть $f_1(x), \dots, f_m(x)$ — заданный набор ортонормированных функций из $L^2(X, \mu)$.

Функцию регрессии $\eta(x)$ будем приближать линейной комбинацией $\sum_{i=1}^m \theta_i f_i(x)$ с помощью метода наименьших квадратов, определяя θ_i путем решения задачи минимизации по $\theta_1, \dots, \theta_m$ суммы квадратов отклонений

$$\sum_{j=1}^N \left(y(x_j) - \sum_{i=1}^m \theta_i f_i(x_j) \right)^2. \quad (4.26)$$

где $y(x_j) = \eta(x_j) + e(x_j)$ — результаты измерений в точках x_j ($j = 1, \dots, N$), $E e(x) = 0$. Полученные таким образом значения $\hat{\theta}_i$ будут отличаться от тех значений θ_i , при которых линейная комбинация $\sum_{i=1}^m \theta_i f_i(x)$ доставляет минимум расстоянию

$$\int_X \left[\eta(x) - \sum_{i=1}^m \theta_i f_i(x) \right]^2 \mu(dx). \quad (4.27)$$

Как известно, этот минимум достигается при $\theta_i = (\eta, f_i)$, где

$$(\eta, f_i) = \int_X \eta(x) f_i(x) \mu(dx)$$

есть коэффициенты Фурье функции η по системе функций $f_1, \dots, f_m$.

Если $p(x_1, \dots, x_N)$ — плотность совместного распределения точек $x_1, \dots, x_N$, а E_p — символ математического ожидания по мере, определяемой этой плотностью, то условия несмещенности плана p записываются в виде

$$E \hat{\theta}_i = (\eta, f_i), \quad i = 1, \dots, m, \quad (4.28)$$

где $E = E_p$ — символ полного математического ожидания (по мере выборочного пространства эксперимента и мере, определяемой плотностью p).

Записывая систему нормальных уравнений для задачи минимизации (4.26) в виде

$$\sum_{i=1}^m \hat{\theta}_i \sum_{j=1}^N f_i(x_j) f_k(x_j) = \sum_{j=1}^N f_k(x_j) y(x_j), \quad k = 1, \dots, m, \quad (4.29)$$

получаем по формуле Крамера

$$\theta_j = \frac{\det \left[\| [f_1, f_1], \dots, [f_{j-1}, f_{j-1}], [f_j, y], [f_{j+1}, f_{j+1}], \dots, [f_m, f_m] \|_{L^2}^m \right]}{\det \left[\| [f_1, f_1], \dots, [f_m, f_m] \|_{L^2}^m \right]} \quad (4.30)$$

где $j = 1, \dots, m$,

$$[f, \psi] = \sum_{i=1}^N f(x_i) \psi(x_i)$$

для любых $f, \psi \in L^2(X, \mu)$.

Отсюда видно, что θ_j можно представить в виде

$$\theta_j = \sum_{i=1}^N y(x_i) A_i^{(j)}(x_1, \dots, x_N),$$

а условия несмещенности (4.28) переписать следующим образом:

$$\begin{aligned} E \hat{\theta}_j &= \sum_{i=1}^N E [y(x_i) A_i^{(j)}(x_1, \dots, x_N)] = \\ &= \sum_{i=1}^N E [\eta(x_i) A_i^{(j)}(x_1, \dots, x_N)] = (\eta, f_j). \end{aligned}$$

Предполагая точки $x_1, \dots, x_N$ равноправными, что естественно, полагаем функцию ρ симметричной относительно своих аргументов. Это приводит к условиям несмещенности в виде

$$\begin{aligned} N \int_X \dots \int_X \eta(x_1) A_1^{(j)}(x_1, \dots, x_N) \rho(x_1, \dots, x_N) \mu(dx_1) \dots \mu(dx_N) = \\ = \int \eta(x) f_j(x) \mu(dx). \end{aligned}$$

Если мы предполагаем, что эти условия выполнены для любой функции $\eta \in L^2(X, \mu)$, то при всех $j = 1, \dots, m$ имеем

$$N \int_X \dots \int_X A_1^{(j)}(x_1, \dots, x_N) \rho(x_1, \dots, x_N) \mu(dx_2) \dots \mu(dx_N) = f_j(x_1). \quad (4.31)$$

Пример 4.2. Пусть $X = [0, 1]$, $f_1(x) = 1$, $f_2(x) = 2\sqrt{3}(x - 1/2)$, $N = m = 2$. В данном случае решение системы нормальных уравнений (4.29) имеет вид

$$\hat{\theta}_1 = \frac{y_1(x_2 - 1/2) - y_2(x_1 - 1/2)}{x_2 - x_1}, \quad \hat{\theta}_2 = (y_2 - y_1)/(x_2 - x_1).$$

Отсюда получаем

$$A_1^{(1)}(x_1, x_2) = (x_2 - 1/2)/(x_2 - x_1), \quad A_1^{(2)}(x_1, x_2) = 1/(x_2 - x_1).$$

Если положить $\rho(x_1, x_2) = 6(x_1 - x_2)^2$, то условия несмещенности (4.31) выполняются. Это свидетельствует о существовании в данном случае несмещенного плана эксперимента.

Сравним построенный рандомизованный план с D -оптимальным планом, который сосредоточен на концах интервала с равными весами. С одной стороны, если функция $\eta(x)$ не является линейной, то применение D -оптимального плана может привести к сколь угодно грубым ошибкам. План, определяемый плотностью $6(x_1 - x_2)^2$, «разбрасывает» точки по отрезку $[0, 1]$ случайно, отдавая предпочтение более удаленным парам точек. Получаемые с помощью этого плана оценки функции регрессии будут иметь своим средним отрезок ряда Фурье функции $\eta(x)$, что гарантирует нас в метрике $L^2[0, 1]$ от грубых просчетов. С другой стороны, при использовании плана $6(x_1 - x_2)^2$ имеем

$$D[\hat{\theta}_1 + \hat{\theta}_2 f_2(x)] = [1 + 12(x - 1/2)^2 + 2\sigma^2] \left[\int_0^1 \eta^2(x) dx - \sum_{i=1}^2 (\eta, f_i)^2 \right]$$

(проверьте!). По сравнению с D -оптимальным планом это не очень хороший результат. Впрочем, можно показать, что существуют более «хорошие» несмещенные планы.

Несмещенные планы, аналогичные плану из предыдущего примера, указаны в формулировке следующей теоремы.

Теорема 4.1. Функция

$$p(x_1, \dots, x_N) = \frac{(N-m)!}{N!} \det \| [f_k, f_i] \|$$

является плотностью по отношению к мере

$$\mu^N(dx_1, \dots, dx_N) = \mu(dx_1) \dots \mu(dx_N)$$

и для нее выполнены условия несмещенности (4.28).

Доказательство. Из ортонормированности функций $f_i(x)$ следует, что матрица $\| (f_i, f_k) \|_{i,k=1}^m$ является единичной. Поэтому выполняется равенство

$$\int \mu^N(dx_1, \dots, dx_N) \frac{(N-m)!}{N!} \det \| [f_k, f_i] \| = 1.$$

Справедливость условий (4.28) легко следует из (4.30) и теоремы 1.21 приложения 1. Теорема доказана.

Из утверждения доказанной теоремы вытекает, что несмещенные планы эксперимента существуют при весьма широких предположениях и дают еще одну характеризацию определителя информационной матрицы: нормированный определитель информационной матрицы является несмещенным планом эксперимента.

Пусть $H(x_1, \dots, x_N)$ — функция, определяющая стоимость измерений в совокупности точек $x_1, \dots, x_N$. Тогда оптимальным несмещенным планом можно назвать плотность

$$p^* = \arg \min_{p \in \mathcal{P}_0} E_p H(x_1, \dots, x_N), \quad (4.32)$$

где $\mathcal{P}_0$ — множество симметричных плотностей распределений p , для которых выполнены условия несмещенности (4.31).

Экстремальная задача (4.32) является задачей линейного программирования (бесконечномерного, если мера μ не сосредоточена в конечном числе точек). Решение такого рода задач линейного программирования — дело трудное. В настоящее время не получен даже для простейших случаев несмещенный аналог D -оптимальных планов, т. е. планов (4.32) для

$$H(x_1, \dots, x_N) = 1/\det \|f_k, f_l\|.$$

Однако не являющиеся оптимальными несмещенные планы, указанные теоремой 4.1, в отдельных случаях могут использоваться на практике. Они особенно удобны при проведении имитационного эксперимента (см. гл. 8).

4.5. Оптимальный выбор рандомизованного плана эксперимента при непараметрическом оценивании регрессии оценками ядерного типа. Пусть $X \subset \mathbb{R}^k$; $K(x, z)$ — функция, заданная на множестве $X \times X$ и близкая к δ -функции $\delta(x - z)$; $x_1, \dots, x_N$ — независимые реализации случайного вектора в X , имеющего плотность распределения $q(x)$; $y(x_1), \dots, y(x_N)$ — результаты измерений в точках $x_1, \dots, x_N$. Положим

$$\hat{\eta}(x) = \frac{1}{N} \sum_{j=1}^N y(x_j) K(x, x_j)/q(x_j). \quad (4.33)$$

Оценка (4.33) для функции регрессии $\eta(x)$ является одной из наиболее распространенных непараметрических оценок и называется оценкой ядерного типа (при этом функция $K(x, z)$ называется ядром).

Лемма 4.2. Пусть выполнено (4.3) и плотность распределения $q(x)$ положительна во всех точках $x \in X$. Тогда для оценки (4.33) имеют место выражения

$$E\hat{\eta}(x) = \int_X \eta(z) K(x, z) dz, \quad (4.34)$$

$$D\hat{\eta}(x) = N^{-1} \left\{ \int_X \frac{\eta^2(z) + \sigma^2(z)}{q(z)} K^2(x, z) dz - \left[\int_X \eta(z) K(x, z) dz \right]^2 \right\}. \quad (4.35)$$

Доказательство. Имеем

$$\begin{aligned} E\hat{\eta}(x) &= E \frac{1}{N} \sum_{j=1}^N y(x_j) K(x, x_j)/q(x_j) = \\ &= E_q \frac{1}{N} \sum_{j=1}^N \eta(x_j) K(x, x_j)/q(x_j) = \\ &= E_q \eta(x_1) K(x, x_1)/q(x_1) = \int_X \eta(z) K(x, z) dz \end{aligned}$$

т. е. получили формулу (4.34). Далее

$$D\hat{\eta}(x) = E\hat{\eta}^2(x) - [E\hat{\eta}(x)]^2$$

Первый член правой части представим в виде

$$\begin{aligned} E\hat{\eta}^2(x) &= EN^{-2} \left\{ \sum_{j=1}^N y(x_j) K(x, x_j)/q(x_j) \right\}^2 = \\ &= N^{-1} \left\{ \sum_{j=1}^N E[y(x_j) K(x, x_j)/q(x_j)]^2 + \right. \\ &\quad \left. + \sum_{\substack{j, l=1 \\ j \neq l}}^N E y(x_j) y(x_l) K(x, x_j) K(x, x_l)/[q(x_j) q(x_l)] \right\} = \\ &= N^{-2} \{ NE[y(x_1) K(x, x_1)/q(x_1)]^2 + \\ &\quad + N(N-1) E y(x_1) y(x_2) K(x, x_1) K(x, x_2)/[q(x_1) q(x_2)] \} = \\ &= N^{-1} \{ E_q [\eta^2(x_1) + \sigma^2(x_1)] K^2(x, x_1)/q^2(x_1) + \\ &\quad + (1 - N^{-1}) [E_q \eta(x_1) K(x, x_1)/q(x_1)] E_q \eta(x_2) K(x, x_2)/q(x_2) \} = \\ &= N^{-1} \left\{ \int_X \frac{\eta^2(x_1) + \sigma^2(x_1)}{q(x_1)} K^2(x, x_1) dx_1 + (1 - N^{-1}) [E\hat{\eta}(x)]^2 \right\}. \end{aligned}$$

Подставляя полученное выражение в формулу для дисперсии, получаем (4.35). Лемма доказана.

Из (4.34) вытекает, что если функция $K(x, z)$ близка к δ -функции $\delta(x - z)$ и функция η удовлетворяет условию Липшица на множестве X , то математическое ожидание оценки (4.33) равномерно близко к значениям $\eta(x)$ при всех $x \in X$. Из (4.9) и (4.35) следует, что при $N \rightarrow \infty$ квадрат случайной погрешности V оценки (4.33) убывает со скоростью N^{-1} . Из леммы 4.2 следует, кроме того, что способ оценивания (4.33) полностью определяет систематическую погрешность, и, следовательно, естественным критерием оптимального выбора плотности распределения является случайная погрешность. В следующей теореме показано, что задача минимизации случайной погрешности по плотности q легко решается.

Теорема 4.2. Пусть имеется схема измерений (4.3), число измерений N , функции $\eta(x)$, $\sigma^2(x)$, $K(x, z)$ фиксированы, непараметрической оценкой функции регрессии $\eta(x)$ является (4.33)

$$k^2(z) = \int_X K^2(x, z) w(x) dx \leq C < \infty, \quad \sigma^2(x) \geq \sigma_0 > 0,$$

а множество допустимых плотностей $q(x)$ состоит из плотностей, всюду положительных на X .

Тогда минимум квадрата суммарной погрешности $I_{\max}$ достигается на плотности

$$q^*(x) = [\eta^2(x) + \sigma^2(x)]^{1/2} k(x) / \int_X [\eta^2(z) + \sigma^2(z)]^{1/2} k(z) dz, \quad (4.36)$$

а также на плотностях, совпадающих с $q^*(x)$ во всех точках, кроме точек из любого подмножества X меры нуль.

Доказательство. Из (4.9) и (4.35) вытекает, что квадрат случайной погрешности V для оценки (4.33) равен

$$V = \int_X [D\hat{\eta}(x)] w(x) dx = \frac{1}{N} \left[\int_X \frac{v^2(z)}{q(z)} dz - \iint_X \eta(z) K(x, z) dz \right]^2 w(x) dx,$$

$$\text{где } v^2(z) = [\eta^2(z) + \sigma^2(z)] \int_X K^2(x, z) w(x) dx.$$

Как указывалось выше, систематическая погрешность не зависит от выбора плотности q , всюду положительной на X . Используя неравенство Коши — Буняковского, убедимся в минимальности квадрата случайной погрешности V для плотности (4.36). Имеем

$$\begin{aligned} NV + \int_X \left[\int_X \eta(z) K(x, z) dz \right]^2 w(x) dx &= \\ &= \int_X \frac{v^2(z)}{q(z)} dz \int_X q(z) dz \geq \left[\int_X \left[\frac{v^2(z)}{q(z)} \right]^{1/2} [q(z)]^{1/2} dz \right]^2 = \\ &= \left[\int_X v(z) dz \right]^2 = \int_X \frac{v^2(z)}{q^*(z)} dz. \end{aligned}$$

Теорема доказана.

Так как функция регрессии $\eta(x)$ априори неизвестна, то неизвестна и оптимальная плотность (4.36) (т. е. оптимальный рандомизованный план эксперимента). Поэтому в практических расчетах следует использовать стандартный последовательный подход. В данном случае суть последовательного подхода состоит в уточнении плотности распределения $q^*(x)$ по мере проведения измерений и оценивания функций $\eta(x)$ и $\sigma^2(x)$.

4.6. Планирование эксперимента при проекционном оценивании функции регрессии. Проекционное оценивание является одним из наиболее распространенных способов непараметрического оценивания регрессии.

При проекционном оценивании регрессии предполагается, что имеется такая возрастающая последовательность $\{L_m\}$ линейных m -мерных пространств, что

$$\mathcal{F} \subset \lim_{m \rightarrow \infty} L_m. \quad (4.37)$$

а число измерений N либо растет, либо постоянно, но достаточно велико. Задача проекционного оценивания состоит в таком выборе размерностей $m = m(N)$ пространств L_m , последовательности заданных до начала измерений точек проведения измерений $x_1, \dots, x_N$ (т. е. плана эксперимента) и метода непараметрического оценивания функции регрессии в предположении $\eta \in L_m$, чтобы получаемая оценка $\hat{\eta}_N(x)$ наилучшим в некотором смысле образом приближала истинную функцию регрессии $\eta \in \mathcal{F}$.

Проекционная оценка функции $\eta(x)$ записывается в виде

$$\hat{\eta}_N(x) = \sum_{i=1}^m \hat{\theta}_i f_i(x) \quad (4.38)$$

где $f_1(x), \dots, f_m(x)$ — базис пространства L_m ; $\hat{\theta}_1, \dots, \hat{\theta}_m$ — оценки параметров линейной функции регрессии

$$\sum_{i=1}^m \theta_i f_i(x) \quad (4.39)$$

построенные по точкам $x_1, \dots, x_N$ и результатам измерений $y(x_1), \dots, y(x_N)$. Как обычно, предполагается, что погрешность оценки функции регрессии определяется по формуле (4.24), что относительно результатов измерений выполнено (4.3) и что метод оценивания является линейным (в данном случае это означает, что статистики $\hat{\theta}_i$ являются линейными функциями от $y(x_1), \dots, y(x_N)$). Будем предполагать также, что $\sigma^2(x) = \sigma^2 > 0$ и что функции $f_1, f_2, \dots$ ортонормированы с весом $w(x)$, т. е. выполнено

$$\int_X f_i(x) f_j(x) w(x) dx = \begin{cases} 1, & \text{если } i = j, \\ 0, & \text{если } i \neq j. \end{cases}$$

При фиксированном m минимально возможное значение квадрата систематической погрешности $B_{\text{сист}}$ равно

$$\mathcal{E}^2(\mathcal{F}, L_m) = \sup_{\eta \in \mathcal{F}} \inf_{\eta_m \in L_m} \int_X (\eta(x) - \eta_m(x))^2 w(x) dx \quad (4.40)$$

Обозначим

$$\theta_i(\eta) = \int_X \eta(x) f_i(x) w(x) dx, \quad i = 1, 2, \dots$$

коэффициенты Фурье функции η по базисным функциям $f_1, f_2, \dots$. Тогда (4.40) может быть переписана в виде

$$\mathcal{E}^2(\mathcal{F}, L_m) = \sup_{\eta \in \mathcal{F}} \sum_{i=m+1}^{\infty} \theta_i^2(\eta). \quad (4.41)$$

Получим теперь нижнюю границу для квадрата случайной погрешности V .

Теорема 4.3. Пусть имеется схема измерений (4.3), $\sigma^2(x) = \sigma^2 > 0$, функции $f_1, f_2, \dots$ ортонормированы с весом $w(x)$ и равномерно ограничены по модулю константой K , статистики $\hat{\theta}_i$ являются линейными, а при фиксированных $m, N, x_1, \dots, x_N$ и $\eta \in L_m$ — несмещенными оценками коэффициентов Фурье $\theta_i(\eta)$ ($i = 1, \dots, m$). Тогда для квадрата случайной погрешности проекционной оценки (4.38) выполняется неравенство

$$V \geq (\sigma^2/K^2) (m/N) \quad (4.42)$$

Доказательство. Пусть m и N фиксированы. Положим

$$f(x) = (f_1(x), \dots, f_m(x))^T, \quad \theta = (\theta_1(\eta), \dots, \theta_m(\eta))^T, \\ \hat{\theta} = (\hat{\theta}_1, \dots, \hat{\theta}_m)$$

В силу (4.9) и ортонормированности функций $f_1, \dots, f_m$ имеем

$$V = \int_X [D \hat{\eta}_N(x)] w(x) dx = \int_X [D(\hat{\theta}^T f(x))] w(x) dx = \\ = \int_X f^T(x) D \hat{\theta} f(x) w(x) dx = \text{tr } D \hat{\theta} \int_X f(x) f^T(x) w(x) dx = \text{tr } D \hat{\theta},$$

где $D \hat{\theta}$ — дисперсионная матрица оценок $\hat{\theta}$.

Оценки $\hat{\theta}$ параметров θ — линейные несмещенные в предположении $\eta \in L_m$. Согласно теореме Гаусса — Маркова (теорема 1.1 гл. I), величина $\text{tr } D\hat{\theta}$ минимальна для случая, когда оценками $\hat{\theta}$ являются оценки МНК. Поэтому далее будем считать, что оценки $\hat{\theta}$ являются МНК-оценками. Для этих оценок

$$D\hat{\theta} = (\sigma^2/N) [M(\xi_N)]^{-1}, \quad (4.43)$$

где

$$M(\xi_N) = \frac{1}{N} \sum_{j=1}^N f(x_j) f^T(x_j)$$

есть нормированная информационная матрица рассматриваемого плана эксперимента, записанного в виде

$$\xi_N = \left\{ \begin{matrix} x_1, & \dots, & x_N \\ 1/N, & \dots, & 1/N \end{matrix} \right\}. \quad (4.44)$$

Пусть теперь ξ — произвольный невырожденный план эксперимента на X ,

$$M(\xi) = \int_X f(x) f^T(x) \xi(dx)$$

есть информационная матрица этого плана. В частности, план ξ может иметь вид (4.44). Применяя к матрице $A = [M(\xi)]^{-1}$ теорему 1.19 из приложения 1, получаем

$$\text{tr } [M(\xi)]^{-1} \geq m^2 / \text{tr } M(\xi) = m^2 / \sum_{i=1}^m \int_X f_i^2(x) \xi(dx) \geq m^2 (mK^2)^{-1} = m/K^2.$$

Учитывая (4.43), получаем (4.42). Теорема доказана.

Процедуру проекционного оценивания функции регрессии будем называть *асимптотически оптимальной*, если при $N \rightarrow \infty$ порядок убывания квадрата суммарной погрешности (4.24) проекционной оценки (4.38) достигает наибольшего значения.

Важную роль при определении максимально возможного порядка убывания квадрата суммарной погрешности $J_{\max}$ играет следующее следствие теоремы 4.3.

Следствие 4.1. Пусть выполнены условия теоремы 4.3 и $N \rightarrow \infty$, $m = m(N) \rightarrow \infty$, $m/N \rightarrow \theta$.

Тогда порядок убывания квадрата суммарной погрешности (4.24) проекционной оценки (4.38) не может быть меньше

$$\mathcal{E}^2(\mathcal{F}, L_m) + m/N. \quad (4.45)$$

Утверждение следствия вытекает из того, что константа σ^2/K^2 в правой части неравенства (4.42) не зависит от m и N .

Опишем одну из простейших процедур проекционного оценивания, для которой при достаточно широких предположениях относительно класса функций $\mathcal{F}$ квадрат суммарной погрешности имеет порядок убывания (4.45).

Пусть $\psi(x) = 1$ и имеется схема измерений (4.3), точки проведения измерений $x_1, x_2, \dots$ в которой представляют собой независимые реализации случайного вектора, имеющего распределение с плотностью $q(x)$, всюду положительной на множестве X . Положим

$$\hat{\theta}_i = N^{-1} \sum_{j=1}^N y(x_j) f_i(x_j) / q(x_j). \quad (4.46)$$

Статистики (4.46) являются простейшими оценками метода Монте-Карло интегралов $\theta_i(\eta)$, т. е. коэффициентов Фурье функции η .

Если числа m и N фиксированы и статистики $\hat{\theta}_i$ выбраны по формуле (4.46), то проекционная оценка (4.38) функции регрессии $\eta(x)$ является частным случаем оценки ядерного типа (4.33), получающейся при

$$K(x, z) = \sum_{i=1}^m f_i(x) f_i(z).$$

Применяя лемму 4.2 получаем, во-первых, что $\hat{\theta}_i$ являются несмещенными оценками коэффициентов Фурье $\theta_i(\eta)$ (поэтому соответствующий рандомизованный план является несмещенным) и, во-вторых, что дисперсионная матрица оценок $\hat{\theta} = (\hat{\theta}_1, \dots, \hat{\theta}_m)^T$ равна

$$D\hat{\theta} = E(\hat{\theta} - \theta)(\hat{\theta} - \theta)^T = N^{-1} \left[\int_X \frac{\eta^2(x) + \sigma^2}{q(x)} f(x) f^T(x) dx - \theta \theta^T \right]. \quad (4.47)$$

Из несмещенности оценок (4.46) следует, что систематическая погрешность $B_{\max}$ получающейся проекционной оценки (4.38) достигает минимально возможного значения:

$$B_{\max} = \mathcal{E}^2(\mathcal{F}, L_m). \quad (4.48)$$

Используя ортонормированность функций $f_1, f_2, \dots$ и (4.47), получаем

$$\begin{aligned} V = V(q) &= E \int_X (E \hat{\eta}_N(x) - \hat{\eta}_N(x))^2 dx = \\ &= E \int_X \left(\sum_{i=1}^m (\theta_i(\eta) - \hat{\theta}_i) f_i(x) \right)^2 dx = \\ &= \sum_{i,j=1}^m E(\hat{\theta}_i - \theta_i(\eta))(\hat{\theta}_j - \theta_j(\eta)) \int_X f_i(x) f_j(x) dx = \text{tr } D\hat{\theta} = \\ &= N^{-1} \left[\int_X \left[\left((\eta^2(x) + \sigma^2) \sum_{i=1}^m f_i^2(x) \right) / q(x) \right] dx - \sum_{i=1}^m \theta_i^2(\eta) \right]. \end{aligned}$$

Согласно (4.36), минимальное значение квадрата случайной погрешности $V(q)$ достигается на плотности

$$q^*(x) = \left[(\eta^2(x) + \sigma^2) \sum_{i=1}^m f_i^2(x) \right]^{1/2} / \int_X \left[(\eta^2(z) + \sigma^2) \sum_{i=1}^m f_i^2(z) \right]^{1/2} dz$$

и равно

$$V(q^*) = N^{-1} \left\{ \left[\int_X \left[(\eta^2(x) + \sigma^2) \sum_{i=1}^m f_i^2(x) \right]^{1/2} dx \right]^2 - \sum_{i=1}^m \theta_i^2(\eta) \right\}. \quad (4.49)$$

Если множество X ограничено, функции $f_1, f_2, \dots$ равномерно ограничены и имеет место

$$\sup_{x \in X} |\eta(x)| < \infty, \quad (4.50)$$

то из (4.49) получаем, что для некоторой константы $C \geq 1$ при $q = q^*$ выполнено неравенство

$$V(q) \leq Cm/N \quad (4.51)$$

Более того, неравенство (4.51) выполнено не только для оптимальной плотности $q = q^*$, но и для многих других плотностей, например, для равномерной плотности $q_0(x) = \text{const}$. Факт существования такой константы $c > 0$, что $V(q^*) \geq cm/N$, очевидно следует из (4.49). Таким образом, для получающейся при $q = q^*$ или $q = q_0$ процедуры проекционного оценивания получаем, что порядок убывания квадрата суммарной погрешности определяется по формуле (4.45). Асимптотически оптимальный выбор последовательности $m = m(N)$ в рассмотренной процедуре проекционного оценивания получается путем приравнивания порядков убывания $V(q)$ и (4.48), т.е. первого и второго слагаемых в (4.45).

Упражнения.

1. Используя условие несмещенности (4.21), покажите, что если существует квадратурная формула

$$\int_X q(x) w(x) dx \approx \sum_{i=1}^n a_i q(x_i),$$

точная для попарных произведений

$$f_l(x) f_j(x) \quad f_l(x) g_k(x), \quad l, j = 1, \dots, m, \quad k = 1, \dots, l.$$

то в качестве точек $x_1, \dots, x_n$ плана (4.15) можно выбрать узлы квадратурной формулы, а в качестве матрицы C из (4.16) — матрицу $C = F_1^T B$, где $F_1 = [f_1(x_1), \dots, f_m(x_1)]_{n-1}^T$, $B = [b_{ij}]_{n-1}^{n-1}$, $b_{ij} = a_i$, $b_{ij} = 0$ ($i, j = 1, \dots, n$).

2. Пусть $X = [-1, 1]$, $f(x) = (1, x, x^2)$, $g(x) = (x^3)$, $x_1 = -0,7746$, $x_2 = -x_1$, $x_3 = 0$, $a_1 = a_2 = 0,2778$, $a_3 = 0,4444$.

Постройте матрицу C согласно способу, указанному в упражнении 1. Покажите, что выполнено условие несмещенности (4.21). Докажите, что для указанных точек x_1, x_2, x_3 план, минимизирующий квадрат случайной погрешности (4.22), имеет вид

$$\xi^* = \begin{Bmatrix} -0,7746 & 0 & 0,7746 \\ 0,2602 & 0,4796 & 0,2602 \end{Bmatrix}.$$

3. Пусть $X = [0, 1]$, $f_1(x) = 1$, $f_2(x) = \begin{cases} 1, & x \in [0, 1/2), \\ -1, & x \in [1/2, 1] \end{cases}$

а) Пусть $N = 2$. Проверьте, что плотность $p(x_1, x_2) = \frac{1}{2} (f_2(x_2) - f_2(x_1))^2$

является несмещенным планом. Покажите, что при использовании этого плана точки x_1 и x_2 не могут находиться в одной и той же половине интервала $[0, 1]$. Вычислите соответствующую этому плану дисперсию оценки функции регрессии по методу наименьших квадратов.

б) Для $N = 3$ вычислите $\frac{1}{3!} \det \|f_k, f_l\|_{k, l=1}^3$. Проверьте несмещенность полученного плана. Вычислите дисперсию оценки функции регрессии

4. Пусть мера μ сосредоточена в точках $0, 1/2, 1$ с равными весами, $f_1(x) = 1$, $f_2(x) = \sqrt{6}(x - 1/2)$, $m = N = 2$. Воспользовавшись теоремой 4.1, укажите несмещенный план. Проверьте его несмещенность. С какими вероятностями должны выбираться пары точек?

5. Пусть функция η принадлежит конечномерному пространству, порожденному заданными функциями $\psi_1, \dots, \psi_m$ из $L^2(X, \mu)$. Какие упрощения условий несмещенности возможны в этом случае? Рассмотрите подробно случаи $m = 2, m = 3$.

6. Пусть $X = [0, 1]$. $w(x) = 1$, $f_1(x) = 1$, $f_{2k}(x) = \sin 2\pi kx$, $f_{2k+1}(x) = \cos 2\pi kx$ ($k = 1, 2, \dots$). $\mathcal{F} = \{\eta \mid \|\theta_i(\eta)\| \leq Lk^{-\alpha} \text{ для } i = 2k \text{ и } i = 2k+1\}$, где $L < \infty$ и $\alpha > 1$ — некоторые константы.

а) Вычислите систематическую и случайную погрешности при равномерной плотности $q(x)$ и способе проекционного оценивания, определенном оценками (4.46). Выберите $m(N)$ из условия равенства V и $B_{\max}$ и покажите, что порядок сходимости к нулю квадрата суммарной погрешности равен $N^{-1+1/(2\alpha)}$.

б) Пусть $x_1, \dots, x_N$ — равноотстоящие точки интервала $[0, 1]$: $x_j = j/N$ ($j = 0, 1, \dots, N-1$), и пусть методом оценивания коэффициентов Фурье $\theta_i(\eta)$ является метод наименьших квадратов. Покажите, что систематическая погрешность проекционной оценки та же, что и в случае а), а случайная погрешность меньше. Выведите отсюда, что получающийся способ проекционного оценивания является асимптотически оптимальным

Глава 5

АНАЛИЗ И ПЛАНИРОВАНИЕ ЭКСПЕРИМЕНТА ДЛЯ НЕЛИНЕЙНЫХ РЕГРЕССИОННЫХ МОДЕЛЕЙ

В предыдущих главах рассматривались схемы регрессионного эксперимента, в которых функция регрессии линейна по неизвестным параметрам. Если эти параметры входят в функцию регрессии нелинейно, то методы регрессионного анализа и планирования регрессионного эксперимента отличаются от рассмотренных выше. В данной главе изложены простейшие из этих методов.

§ 1. Нелинейный регрессионный анализ

1.1. Нелинейные регрессионные модели. Рассмотрим следующую схему, которая является обобщением линейных схем, изученных в предыдущих главах. Ее также называют *схемой регрессионного эксперимента*. Пусть в точках x_j ($j = 1, \dots, N$) из множества планирования X наблюдаются значения случайных величин $y_j = y(x_j)$, представимых в виде

$$y_j = \eta(x_j, \theta) + \varepsilon_j, \quad j = 1, \dots, N, \quad (1.1)$$

где $\eta(x, \theta)$ — функция на $X \times \Omega$, называемая *функцией регрессии* и заданная с точностью до неизвестных параметров $\theta = (\theta_1, \dots, \theta_m)^T$ из некоторого параметрического множества $\Omega \subset R^m$, ε_j — случайные ошибки измерений. Как обычно, будем предполагать, что ошибки измерений ε_j центрированы ($E\varepsilon_j = 0$) и некоррелированы ($E\varepsilon_i \varepsilon_j = 0$ при $j \neq i$), а измерения равноточны ($D\varepsilon_j = E\varepsilon_j^2 = \sigma^2$), причем значение σ^2 может быть неизвестным.

В случае, когда функция регрессии $\eta(x, \theta)$ линейно зависит от неизвестных параметров θ , т. е.

$$\eta(x, \theta) = \theta^T f(x), \quad (1.2)$$

где $f(x) = (f_1(x), \dots, f_m(x))^T$ — вектор известных (базисных) функций, схема регрессионного эксперимента является схемой линейной регрессии.

Если функция регрессии $\eta(x, \theta)$ нелинейна по параметрам θ , то она называется *нелинейной*, а схема регрессионного эксперимента (1.1) называется *нелинейной регрессионной моделью*.

Несмотря на то, что задача оценивания неизвестных параметров нелинейных регрессионных моделей значительно сложнее аналогичной задачи для линейных моделей, нелинейные регрессионные задачи на практике встречаются довольно часто. Это происходит по следующим причинам: а) нелинейность исходит из сущности явления, для описания которого предназначена модель; б) дополнительная информация об истинном характере зависимости часто позволяет выбрать достаточно точную (т. е. имеющую малую систематическую ошибку) нелинейную модель с числом параметров, значительно меньшим, чем для аналогичной линейной модели; в) линеаризация модели часто приносит значительно больше потерь, чем выгод (см. п. 1.2). Примерами нелинейных функций регрессии, широко распространенных в практике экспериментальных исследований (в частности, в физике, химии, биологии), могут служить следующие.

Пример 1.1. Функция регрессии

$$\eta(x, \theta) = \sum_{i=1}^m \theta_i \exp \{-\theta_{i+m} x\} \quad (1.3)$$

возникает в результате решения систем линейных дифференциальных уравнений и широко применяется при исследовании, в частности, кинетических задач.

Пример 1.2. Функции регрессии

$$\eta(x, \theta) = \sum_{i=1}^m \theta_i / [(x - \theta_{i+m})^2 + \theta_{i+2m}]$$

описывают резонансные явления и используются при обработке результатов спектрального анализа. В этих же задачах используются и функции регрессии, имеющие вид

$$\eta(x, \theta) = \sum_{i=1}^m \theta_i \exp \{-\theta_{i+m} (x - \theta_{i+2m})^2\}.$$

Пример 1.3. Функции регрессии

$$\eta(x, \theta) = \arctan g[(x_1 - \theta_1)/(x_2 - \theta_2)], \quad x = (x_1, x_2)^T,$$

используются в задачах слежения за движущимися объектами.

Пример 1.4. Если известно, что функция регрессии унимодальная (т. е. имеет один локальный максимум), достаточно гладкая и существует такое θ_1 , что $\eta(x) = 0$ при $x \leq \theta_1$ и $\lim_{x \rightarrow \infty} \eta(x) = 0$, но неизвестен механизм процессов, приводящих к изучаемой зависимости, то в качестве функции регрессии можно выбирать

$$\eta(x, \theta) = \theta_2 (x - \theta_1)^{\theta_3} \exp \{-\theta_1 (x - \theta_1)^{\theta_3}\}, \\ x \geq \theta_1, \quad \theta_i \geq 0, \quad i = 2, 3, 4, 5.$$

Метод наименьших квадратов — один из наиболее распространенных методов оценивания параметров нелинейных регрессионных моделей.

Оценкой МНК неизвестных параметров θ в схеме регрессионного эксперимента (1.1) называется оценка вида

$$\hat{\theta} = \hat{\theta}_N = \arg \min_{\theta \in \Omega} \sum_{i=1}^N (y_i - \eta(x_i, \theta))^2. \quad (1.4)$$

Ясно, что (1.4) является обобщением МНК-оценки (1.1) из гл. 1 на случай нелинейной функции регрессии.

Оценивание по МНК не единственно возможный способ оценивания неизвестных параметров. Пусть, например, ошибки ε_i взаимно независимы и имеют одинаковое известное с точностью до σ распределение вероятностей с плотностью p_σ . Тогда для оценивания θ и σ можно использовать оценки максимального правдоподобия

$$(\hat{\theta}, \hat{\sigma}) = \arg \max_{\substack{\theta \in \Omega, \\ \sigma > 0}} L(\theta, \sigma, y_1, \dots, y_N),$$

где функция правдоподобия L равна

$$L(\theta, \sigma, y_1, \dots, y_N) = \prod_{i=1}^N p_\sigma(y_i - \eta(x_i, \theta)).$$

Если $p_\sigma(t)$ — плотность нормального распределения

$$p_\sigma(t) = \frac{1}{\sqrt{2\pi}\sigma} e^{-t^2/(2\sigma^2)},$$

то, как легко показать (покажите!), оценки максимального правдоподобия параметров θ совпадают с МНК-оценками. Ситуация аналогична ситуации для линейного случая, рассмотренной в § 1 гл. 1.

1.2. Возможные последствия линеаризации модели. Предположим, что по результатам измерений (1.1) требуется оценить параметры нелинейной функции регрессии $\eta(x, \theta)$, которая после некоторого преобразования $\Phi: \mathbb{R} \rightarrow \mathbb{R}$ становится линейной по параметрам. В этом случае для оценивания параметров θ регрессии $\eta(x, \theta)$ можно использовать два подхода, основанных на принципе наименьших квадратов. Первый подход — прямое построение МНК-оценок (1.4) для схемы измерений (1.1), т. е. минимизация по θ функции

$$Q(\theta) = \sum_{i=1}^N (y_i - \eta(x_i, \theta))^2. \quad (1.5)$$

Второй подход основан на линеаризации схемы измерений (1.1), т. е. на рассмотрении вместо нее схемы

$$z_i = \Phi(y_i) = \Phi(\eta(x_i, \theta)) + \delta_i, \quad i = 1, \dots, N. \quad (1.6)$$

с последующей оценкой по МНК получившейся линейной регрессионной модели. Здесь $\delta_i = \Phi(y_i) - \Phi(\eta(x_i))$ — некоррелированные случайные величины, распределение которых зависит от распределения ε_i , от Φ и η .

На очень простом, но типичном примере покажем, что, несмотря на значительный выигрыш в простоте, проигрыш в точности получаемых во втором подходе оценок может быть столь существенным, что делает линеаризацию модели нежелательной. В качестве исходной рассмотрим схему измерений

$$y_i = \sqrt{x_i + \theta} + \varepsilon_i, \quad i = 1, \dots, N, \quad (1.7)$$

где случайные ошибки ε_i независимы и нормально распределены с нулевым средним и дисперсией σ^2 , число измерений N достаточно велико, а точки измерений x_i ($x_i \geq -\theta$) выбраны согласно некоторому непрерывному плану эксперимента $\xi(dx)$ на $[-\theta, \infty)$ с плотностью $p(x)$.

В данном случае функция регрессии имеет вид $\eta(x, \theta) = \sqrt{x + \theta}$. После возведения в квадрат получаем

$$\eta^2(x, \theta) = x + \theta$$

Эта функция линейна как функция от неизвестного параметра. Величины z_i , δ_i в данном случае равны

$$z_i = y_i^2 = (\sqrt{x_i + \theta} + \varepsilon_i)^2, \\ \delta_i = (\sqrt{x_i + \theta} + \varepsilon_i)^2 - (x_i + \theta) = \varepsilon_i^2 + 2\varepsilon_i \sqrt{x_i + \theta}.$$

Оценку, полученную по формуле (1.4), будем далее называть *нелинейной* и обозначать $\hat{\theta}_N(n)$:

$$\hat{\theta}_N(n) = \arg \min_{\theta} \sum_{i=1}^N (y_i - \sqrt{x_i + \theta})^2. \quad (1.8)$$

МНК-оценку, построенную по схеме измерений

$$z_i = x_i + \theta + \delta_i, \quad i = 1, \dots, N, \quad (1.9)$$

в предположении центрированности, некоррелированности и равноточности случайных величин δ_i , назовем *линеаризованной* и будем обозначать $\hat{\theta}_N(l)$. Эта оценка определяется по формуле

$$\hat{\theta}_N(l) = \arg \min_{\theta} \sum_{i=1}^N (z_i - (\theta + x_i))^2$$

и равна, как нетрудно проверить (проверьте!),

$$\hat{\theta}_N(l) = \frac{1}{N} \sum_{i=1}^N (z_i - x_i).$$

Вычислим асимптотические (при $N \rightarrow \infty$) характеристики оценок $\hat{\theta}(n)$ и $\hat{\theta}(l)$. Сначала воспользуемся общей теоремой 1.1 из п. 1.3 об асимптотическом распределении нелинейных оценок наименьших квадратов. Согласно этой теореме, в случае одного параметра ($m = 1$) и проведения измерений по непрерывному плану с плотностью $p(x)$ асимптотическое (при $N \rightarrow \infty$) распределение МНК-оценок является нормальным с нулевым средним и дисперсией

$$\sigma_N^2 = \frac{\sigma^2}{N} \int_{-\infty}^{\infty} \left(\frac{\partial \eta(x, \theta)}{\partial \theta} \right)^2 p(x) dx.$$

В рассматриваемом случае

$$\frac{\partial \eta(x, \theta)}{\partial \theta} = \frac{\partial \sqrt{x + \theta}}{\partial \theta} = \frac{1}{2\sqrt{x + \theta}},$$

поэтому

$$\sigma_N^2 = \frac{\sigma^2}{4N} \int_{-\infty}^{\infty} \frac{1}{x + \theta} p(x) dx. \quad (1.10)$$