

АКАДЕМИЯ НАУК БЕЛОРУССКОЙ ССР
ИНСТИТУТ ТЕХНИЧЕСКОЙ КИБЕРНЕТИКИ

А. Д. ЗАКРЕВСКИЙ

ЛОГИКА распознавания

МИНСК
«НАУКА И ТЕХНИКА»
1988

ББК 32.81
3-20
УДК 681.327.12

Рецензенты:
д-р техн. наук П. М. Чеголин,
канд. техн. наук Ю. Н. Печерский

Закревский А. Д.

3-20 **Логика распознавания.**— Мн.: Наука и техника,
1988.— 118 с.

ISBN 5-343-00404-0.

Сложное поведение кибернетических систем складывается из цепочек «распознавание ситуации—принятие решения—действие». Книга посвящена первому из этих звеньев. В популярной форме излагается подход к распознаванию, основанный на построении «модели мира» — логического пространства взаимосвязанных признаков и разделения процесса распознавания на два этапа: индуктивный и дедуктивный.

Предназначена для читателей, интересующихся современными проблемами кибернетики.

1502000000—076
3 ————— 148—88
М316(03)—88

ББК 32.81

ISBN 5-343-00404-0.

© Издательство
«Наука и техника», 1988.

Как только не называют наш век — и веком электричества, и веком атомной энергии, и веком пластмасс, и даже веком очередей... Однако наиболее емким и обоснованным выглядит термин «информационный век».

В самом деле, по темпам развития и по широте применения электронные вычислительные и управляющие устройства — а именно они перерабатывают в наши дни основные потоки информации — оставили далеко позади и энергетические машины, и технологию производства вещественных изделий.

Пользуясь автомобилем, мы передвигаемся быстрее в десятки раз, а с помощью самолета — в добрую сотню раз. Для информационных процессов эффект машинного ускорения гораздо выше: современная ЭВМ считает в миллион раз быстрее, чем человек, а перспективные лабораторные образцы — даже в миллиард раз!

Чтобы решить некоторую задачу на ЭВМ, надо прежде всего свести ее к вычислениям. Причем к вычислениям в расширенном смысле — не только арифметическим, но и логическим. Арифметические вычисления опираются на солидный фундамент вычислительной математики. С логическими вычислениями труднее — вычислительная логика находится в самом начале своего развития, закладываются лишь первые камни в ее основание. Тем не менее логическим вычислениям, логическому программированию в последнее время уделяется все большее внимание, поскольку они позволяют значительно расширить круг решаемых на ЭВМ задач, выйти за пределы традиционного применения ЭВМ.

Что же это за задачи?

Принято говорить, что они относятся к

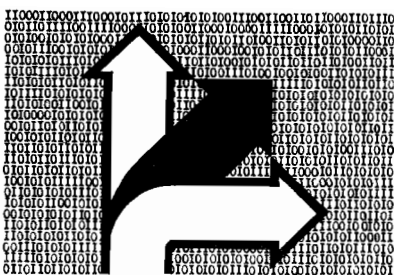
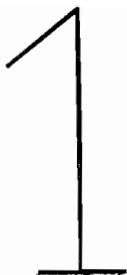
области искусственного интеллекта. А ближайшие практические результаты в этой области связываются с построением экспертных систем (ЭС). Так называются поставленные на ЭВМ программные комплексы, имитирующие поведение экспертов в конкретных предметных областях: медицине, строительном деле, проектировании дорог и т. д.

Как правило, экспертная система содержит базу данных, в которой накапливаются различные факты из предметной области, и базу знаний, где фиксируются некоторые правила, или закономерности, присущие данной области. В ЭС имеются средства общения со специалистами-экспертами, из которых система «выкачивает» знания, и со специалистами-пользователями, которым система помогает в их работе. Хорошо экипированные ЭС обладают средствами индуктивной логики, автоматизирующими переход от данных к знаниям, и практически в любую ЭС встроен механизм логического (дедуктивного) вывода — от накопленных знаний к заключениям в конкретных ситуациях.

Предметные области разнообразны, и их специфика должна учитываться при разработке эффективных ЭС. Однако при всем этом разнообразии практически все ЭС опираются в конечном счете на логику, которая едина, хотя и сильно разветвилась к настоящему времени.

Цель книги — познакомить читателя с логическими основаниями распознающих экспертных систем, ориентированных на относительно простые предметные области, описываемые в терминах булевой алгебры. К счастью, такими оказываются многие практически интересные области.

Изложение ведется на популярном уровне, для чтения книги вполне достаточно знаний в объеме средней школы.



НАДЕЖНОЕ РАСПОЗНАВАНИЕ — ОСНОВА РАЗУМНЫХ ДЕЙСТВИЙ

Поведение живых организмов складывается из элементарных актов, каждый из которых можно разложить на три операции: оценку текущей ситуации, выработку решения и его реализацию. Поведение дискретно. Дискретность сглаживается у массивных животных инерцией тела, но хорошо прослеживается при наблюдениях над насекомыми. Ее трудно заметить, глядя на молодого человека, полного сил и энергии, но она проявляется в движениях старика, поскольку с возрастом нервная система постепенно изнашивается, реакции замедляются, промежутки времени, отводимые под элементарные акты поведения, становятся длительнее.

Поведение иерархично, как и материя. Сложные акты поведения состоят из более простых, которые в свою очередь могут распадаться на еще более мелкие. Дискретность поведения неизбежна. Прежде чем перейти улицу, надо решить, в каком именно месте это сделать, и выбрать подходящий момент времени. Следует проанализировать текущую ситуацию: посмотреть на светофор, если он есть, оценить обстановку на проезжей части улицы. На каждое действие затрачивается какое-то время. Сделав один шаг, надо учесть возможные измене-

ния ситуации, прежде чем сделать второй. Только при этом условии поведение разумно.

В последние годы много говорят и пишут о «разумных машинах», или о машинах с «искусственным интеллектом». Имеются в виду не простейшие автоматические устройства, например открывающие дверь в холле гостиницы при приближении любого человека. Речь идет о более сложных операциях, выполняемых машиной, например, дверь открывается лишь в том случае, когда к ней подходит человек, фотокарточка которого заранее передана машине.

В данном примере наиболее сложной оказывается первая операция акта поведения — оценка ситуации, или опознание подошедшего человека. Именно она является ключевой операцией, определяющей эффективность системы управления дверью.

Простые операции распознавания широко и уже давно используются в технических системах, выполняются автоматически несложными механизмами.

Английский писатель-экономист С. Лилли упоминает о рабочем, который «много лет подряд выполнял одну и ту же операцию по затяжке болта на беспрерывно движущемся конвейере по сборке задних мостов. Он перешел работать на консервный завод, где должен был сортировать черешню, по мере того как она подавалась на конвейере,— черную налево, белую направо. У него были хорошие условия, отличная зарплата. Однако в конце недели он попросил расчет, заявив: «Ответственность — вот что меня удручает. Всегда думай, думай, думай!» *

Этим анекдотичным примером С. Лилли иллюстрирует мысль о том, что однообразная работа отучает от творчества и ответственности. Автоматическое исполнение простых повторяющихся действий превращает человека в автомат, и этот автомат тем тривиальнее, чем проще выполняемая им работа. Создавая условия для развития творческих способностей человека, мы должны освободить его от рутинной работы, поручив ее машинам.

Простые операции распознавания иногда могут выполняться одновременно с действиями, основанными на результатах распознавания, по существу сливаться с

* Лилли С. Автоматизация и социальный прогресс. М., 1958.

ними, как это видно на примере просеивания муки через сито. Но при усложнении процессы распознавания отделяются от последующих действий, в них более четко выступает информационная сущность, они приобретают определенную автономность.

Разумные машины должны уметь распознавать ситуации, в которых им предстоит действовать. А научить их этому можно, лишь располагая методами распознавания, причем достаточно формальными, иначе машина не поймет объяснений. Автоматизации процессов должна предшествовать их формализация!

Не случайно поэтому в новом научном направлении с интригующим названием «Искусственный интеллект» большое внимание уделяется исследованиям в области формальных методов распознавания.

Задачи распознавания разнообразны как по постановке, так и по методам решения. Относительно просты операции распознавания «по определению». Например, млекопитающее определяется как животное, дающее молоко. Следовательно, чтобы утверждать, что некоторый предмет — млекопитающее, достаточно убедиться в том, что, во-первых, это животное, а не какая-либо машина и, во-вторых, что он действительно вырабатывает молоко. Подобным образом строятся правила распознавания и других естественных объектов, подвергнутых классификации: животных, растений, минералов, метеорологических явлений и т. д. Распознавать их помогают справочники-определители, разработанные для отдельных классов объектов и перечисляющие для каждого вида представителей этих классов те свойства, которыми этот вид должен обладать, а также, возможно, те свойства, которые у него должны отсутствовать.

Аналогичный подход широко используется в диагностике, начальное представление о которой можно составить на основе следующей простой модели.

Пусть a, b, c, d, e' — скрытые свойства объектов исследуемого класса (например, конкретные неисправности внутри микроэлектронной схемы или болезни человека), недоступные непосредственному наблюдению, а p, q, r, s — наблюдаемые свойства (например, состояния внешних полюсов схемы или некоторые симптомы болезней). Предположим для простоты, что любой конкретный объект обладает одним и только одним из скрытых свойств. Допустим также, что связь между скрытыми

и наблюдаемыми свойствами представляется следующей таблицей:

	<i>p</i>	<i>q</i>	<i>r</i>	<i>s</i>
<i>a</i>	0	1	0	1
<i>b</i>	1	0	0	0
<i>c</i>	1	1	1	0
<i>d</i>	1	0	0	1
<i>e</i>	0	1	1	1

Например, первая строка говорит о том, что если некоторый объект обладает свойством *a*, то он неизбежно обладает и свойствами *q* и *s*, но не обладает ни свойством *p*, ни свойством *r*. Таким образом, каждая строка определяет соответствующее скрытое свойство у рассматриваемых объектов, выражая его через другие свойства, доступные для наблюдения.

Формулируемая на этой модели задача заключается в распознавании скрытых свойств: наблюдая за свойствами *p*, *q*, *r*, *s*, т. е. устанавливая их наличие или отсутствие у некоторого конкретного объекта, требуется решить, каким из свойств *a*, *b*, *c*, *d*, *e* он обладает. Нетрудно сообразить, что данную задачу можно решить полностью лишь в том случае, если все строки таблицы различны, в противном случае соответствующие скрытые свойства будут неразличимы.

В данном примере такое условие выполняется. Более того, удалив из таблицы второй столбец, соответствующий свойству *q*, мы получим остаток, в котором тоже не будет одинаковых строк. Это значит, что при решении задачи можно обойтись без «измерения» свойства *q*.

Не составляет особого труда построение алгоритма вычисления скрытого свойства в рассматриваемом примере. Можно сначала «измерить» свойства *p*, *r* и *s*, а затем сравнить полученные результаты со строками таблицы и определить скрытое свойство «по совпадению». А можно и сэкономить на измерениях, предложив следующий алгоритм распознавания.

1. Измерить *p*. Если *p*=1 (т. е. свойство *p* налицо), то перейти к п. 2, в противном случае — к п. 4.

2. Измерить *r*. Если *r*=1, то сделать вывод «Объект обладает свойством *c*» и перейти к п. 5. В противном случае — к п. 3.

3. Измерить *s*. Если *s*=1, то сделать вывод «Объект

обладает свойством d » и перейти к п. 5. В противном случае сделать вывод «Объект обладает свойством b » и перейти к п. 5.

4. Измерить r . Если $r=1$, то сделать вывод «Объект обладает свойством e » и перейти к п. 5. В противном случае сделать вывод «Объект обладает свойством a » и перейти к п. 5.

5. Закончить выполнение алгоритма.

К сожалению, большинство практических задач распознавания не укладывается в описанную схему. В более сложных ситуациях связь между наблюдаемыми и вычисляемыми свойствами не столь проста. Более того, она может носить фрагментарный характер, оставляя место различного рода домыслам и гипотезам. Таковы связи, с которыми имеет дело известный подход «обучения на примерах».

Допустим, что Вы впервые в жизни собираетесь в лес за грибами и желаете узнать у специалиста-грибника, как отличать съедобные грибы от несъедобных. Вместо разъяснений он показывает Вам две корзины: в одной собраны съедобные грибы, а во второй — несъедобные. Однако в корзинках лежат далеко не все грибы, которые могут встретиться в лесу! Что делать, если Вам попадется гриб, не похожий на собранные? Ответ на этот вопрос Вы должны найти самостоятельно.

Возникающая здесь задача (если отвлечься от конкретной интерпретации) — одна из наиболее известных в теории распознавания. Решить ее — значит построить правила распознавания, применимые к любым объектам интересующего нас типа (например, к любым грибам). Сложность задачи заключается в том, что, располагая ограниченной информацией, мы должны экстраполировать ее, а сделать это можно по-разному. Поэтому многочисленные методы решения данной задачи часто приводят к различным результатам.

Не будем пытаться дать в предлагаемой вниманию читателя книге обзор всех этих методов. Ограничимся более скромной целью — изложим в популярной форме оригинальный подход к распознаванию, предложенный автором в серии научных публикаций *.

* Закревский А. Д. К формализации полисиллогистики // Логический вывод. М., 1979.

Закревский А. Д. Некоторые комбинаторные задачи искусственного интеллекта // Семиотика и информатика. 1980. Вып. 15.

Воспользуемся традиционным разделением процесса распознавания на два этапа: обучение и собственно распознавание. Первый этап индуктивный, второй — дедуктивный. На первом из них обрабатываются данные многочисленных наблюдений над отдельными представителями исследуемого класса объектов (вспомним о лежащих в корзинках грибах) и на основе полученных результатов строится некоторое решающее правило. На втором этапе описанное правило применяется для распознавания интересующих нас, но непосредственно не измеряемых свойств других объектов этого же класса.

На этапе обучения выявляются некоторые закономерности, присущие исследуемому классу, и совокупность этих закономерностей служит далее «моделью мира». На ее основе решаются задачи распознавания свойств конкретных объектов. Очевидно, что «модели мира», а, проще говоря, формальные модели исследуемых классов объектов будут различаться: для грибов одна модель, для минералов другая, для органических молекул третья и т. д.

Оригинальность подхода заключается в выборе типа, или формы рассматриваемых закономерностей. Предлагается ограничиться поиском закономерностей типа «элементарный запрет». Такой выбор позволяет привлечь для решения задач второго этапа хорошо развитый аппарат алгебры логики. Поскольку проводимые на первом этапе наблюдения, как правило, ограничены, выявляемые закономерности по необходимости носят характер гипотез. Поэтому наряду с четким распознаванием некоторых свойств исследуемых объектов подход предусматривает отказ от распознавания (в силу недостаточности информации, полученной на первом этапе),

Закревский А. Д. К выбору гипотез при поиске закономерностей в булевом пространстве // Докл. АН БССР. 1981. Т. 25, № 3.

Закревский А. Д. Выявление имплицативных закономерностей в булевом пространстве признаков и распознавание образов // Кибернетика. 1982. № 1.

Закревский А. Д., Данг Динь Куанг. К выявлению функциональных закономерностей в булевом пространстве признаков // Докл. АН БССР. 1983. Т. 27, № 3.

Закревский А. Д. Синтез дедуктивной системы распознавания на основе выявления имплицативных закономерностей // Применение методов математической логики. Представление знаний и синтез программ: Тез. IV Всесоюзн. конф. Таллин, 1986.

а также обнаружение противоречий между результатами наблюдений на втором этапе и принятыми ранее гипотезами.

По-видимому, основным понятием при таком подходе является понятие закономерности. Трудно найти его определение в каком-либо словаре. Похоже, что вначале появилось слово «закономерный», т. е. «соответствующий, отвечающий законам», как говорится в Словаре русского языка С. И. Ожегова. А для понятия закона там приведены три значения: 1) связь и взаимозависимость каких-нибудь явлений объективной действительности; 2) постановление государственной власти; 3) общеобязательное правило, то, что признается обязательным. Под закономерностью обычно понимается нечто менее строгое, чем закон, в первом из этих значений. Направленное на широкий круг явлений, оно определено довольно расплывчато. Но в приложении к более ограниченным, но зато более строгим формальным моделям его можно уточнять, делая приемлемым на математическом уровне. Этому и посвящена значительная часть настоящей книги.

Выявление закономерностей в потоке данных — основной путь научного познания окружающей нас действительности. И надо сказать, чрезвычайно трудоемкий, даже если заранее предугадать, какие именно величины связываются искомой закономерностью.

Прекрасным примером этому служит открытие Кеплером одного из знаменитых законов, носящих его имя. Предчувствуя закономерную связь между средними расстояниями планет от Солнца и временем их обращения вокруг него, Кеплер, по свидетельству Лапласа, «долго сравнивал их как с правильными геометрическими телами, так и с интервалами звуковых тонов. Наконец, после 17 лет бесполезных проб, решив сравнить степени расстояний с временами звездных обращений, он нашел, что квадраты этих времен относятся между собой как кубы больших осей орбит. Этот очень важный закон, который ему удалось получить и для системы спутников Юпитера, распространяется на все системы спутников».

Приведенный пример демонстрирует индуктивный путь получения научных результатов. В своих изысканиях Кеплер опирался на результаты многолетних астрономических наблюдений своего учителя Тихо Браге. Сформулировав три закона, Кеплер экстраполировал

данные многочисленных конкретных наблюдений за планетами и уложил их в сжатую математическую форму, позволяющую предсказывать все многообразие возможных положений планет в любые моменты времени.

Опираясь на эти капитальные результаты, Ньютон открыл впоследствии закон всемирного тяготения: сила взаимного притяжения двух материальных частиц (размеры которых настолько малы, что ими можно пренебречь) пропорциональна произведению их масс и обратно пропорциональна квадрату расстояния между ними. Это — пример функциональной связи между входящими в закон четырьмя величинами: зная значение трех величин, можно однозначно определить значение четвертой.

В последующем развитии физики функциональная связь стала традиционной и почти единственной формой, в которую укладывались многочисленные закономерности, выявляемые в научных исследованиях. Между тем функциональная связь — не единственная форма закономерности. Ее можно рассматривать лишь как частный случай более общего типа закономерности — запрета, задающего в пространстве описаний объектов исследуемого класса некоторые компактные области, где заведомо или с высокой степенью достоверности не существует ни одного объекта. Система таких закономерностей образует общую область запрета, знание которой помогает предсказывать особенности конкретных объектов с-частично заданными свойствами.

Закономерности в форме запретов представлены довольно широко. Наиболее очевидно это при рассмотрении правил, регулирующих человеческое поведение — в обществе и при контактах с техникой. Вспомним такие таблички, как «Здесь не курить», «Просьба до 14⁰⁰ не беспокоить», «Уходя, гасите свет». Все они представляют запреты в прямой или слегка завуалированной форме. Эта форма оказалась удобной в биологии, где связи между отдельными свойствами изучаемых объектов часто формулируются в виде импликаций, например типа «если объект обладает свойством a , то он не может обладать свойством b » (другими словами, комбинация свойств a , b оказывается под запретом). Да и в современной физике эта форма используется не так уж редко — вспомним запрет Паули, принцип неопределенности Гейзенберга.

Закономерностям такого рода, выражаемым в виде областей запрета с достаточно простой формой, подчиняется и распределение материи в привычном для нас трехмерном пространстве. В результате проведенных астрономических исследований, направленных на изучение макроструктуры Вселенной, обнаружено существование областей, напоминающих громадные мыльные пузыри размером в миллионы и миллиарды световых лет, внутри которых не существует ни галактик, ни других сгустков материи — они как бы запрещены в этих областях, пронизываемых лишь слабым излучением далеких звезд.

Пространство, о котором мы будем говорить далее, иного рода. Это — искусственное многомерное пространство, число измерений в котором достигает нескольких десятков и даже сотен, что вызывает определенные трудности при его исследовании. Но оно проще в другом отношении: каждая координата может принимать значения из множества, содержащего лишь два элемента, которые обычно обозначаются через 0 и 1. Такое пространство называется булевым, и его точки служат моделями реальных объектов. При этом координатами представляются некоторые свойства объектов: если данное свойство объекту присуще, соответствующая координата имеет значение 1, в противном случае — 0. Разумеется, эти модели ограничены, поскольку реальные объекты всегда обладают массой других свойств, не учитываемых при таком моделировании — соответствующие им координаты в пространство не вводятся. Однако любая модель строится для решения некоторого ограниченного круга задач, и с этой точки зрения многие из свойств реальных объектов оказываются несущественными, что и позволяет исключать их из дальнейшего рассмотрения.

Наука дифференцирована, в каждой ее области изучается свой класс объектов, и при их описании используются свойства, характерные для данного класса.

Исследуемый класс объектов в целом можно представить как множество точек, распределение которых в булевом пространстве признаков (так в теории распознавания принято называть свойства объектов) подчиняется некоторым закономерностям: где-то образуются сгущения, где-то разрежения, а где-то и совершенно пустые области — аналогично распределению звезд во

Вселенной. Но обычно это множество неизвестно — оно слишком велико. В практических ситуациях известна лишь малая часть образующих его точек, представляющая данные некоторых наблюдений и экспериментов. На основе имеющейся ограниченной информации требуется составить представление о классе в целом, а значит, выявить присущие ему закономерности. В этом и заключается индуктивный этап распознавания.

Оригинальной для предлагаемого нами подхода является следующая конкретизация понятия закономерности: под закономерностью понимается связь между признаками, причем достаточно сильная, чтобы ее можно было выявить на индуктивном этапе. А связь между признаками — это запрет на некоторую комбинацию их значений. Чем меньше признаков связано, тем сильнее связь и тем легче обнаружить ее при анализе экспериментальных данных.

Заметим, что все индуктивные выводы представляют собой не что иное, как некоторые гипотезы. Если, к примеру, при экспериментальном исследовании некоторого класса не обнаружены объекты, одновременно обладающие признаками «горячий» и «красный», то это еще не означает, что таких объектов вообще нет. Но чем больше исследовано конкретных объектов, тем больше оснований для выдвижения такой гипотезы. Возникает проблема оценки ее достоверности. Проведенные исследования привели к интересным результатам, излагаемым в последующих главах.

Конечный результат индуктивного этапа распознавания представляется совокупностью выявленных закономерностей. Каждая из них задает в булевом пространстве признаков некоторую запретную область простой формы, так называемый интервал, а все вместе — интегральную запретную область более сложного вида: утверждается, что в этой области объектов не существует.

Обнаруженные закономерности играют роль аксиом в процессе логического вывода на втором этапе распознавания — дедуктивном. Основная идея такого вывода заключается в следующем. Допустим, исследуется некоторый новый объект того же класса. Установлено, что одними признаками он обладает, других не имеет. Относительно остальных признаков ничего не известно — то ли они присущи данному объекту, то ли нет. Непосред-

ственное «измерение» этих признаков в силу каких-то причин невозможно, поэтому целесообразно их «вычислить». Имеющуюся информацию об исследуемом объекте можно представить в форме утверждения о том, что объект расположен в определенном интервале булева пространства признаков. Если интервал пересекается с запретной областью, то область возможного существования объекта уменьшается — в этом и заключается логический вывод. Далее все зависит от величины и, главное, от формы полученной суженной области. При благоприятном стечении обстоятельств можно получить четкие логические выводы о наличии или отсутствии у объекта некоторых признаков из числа первоначально не измеренных. Например, если вычисленная область возможного существования объекта не пересекается с той половиной пространства, которая соответствует значению 1 признака a , то это означает, что данный признак объекту не присущ.

Однако формализовать проверку такого пересечения, перевести ее на алгебраический язык не так-то просто.

Если выявленные закономерности связывают только два признака каждая, то они представляют собой не что иное, как простые суждения, на которых построена силлогистика Аристотеля. Таким образом, силлогистика оказывается исторически первой формальной распознающей системой. В ней рассматриваются простейшие умозаключения: в классической фигуре силлогизма всего три простых суждения, два из которых играют роль посылок, а одно — логического следствия.

Например, в известном силлогизме из двух посылок

Все люди смертны,
Сократ—человек

выводится логическое следствие

Сократ—смертен.

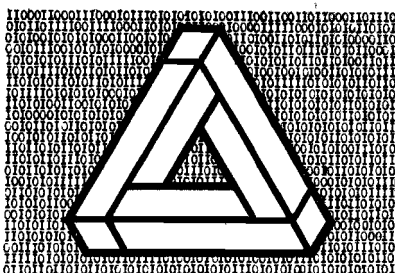
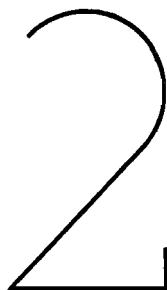
Логический вывод, которому посвящена последняя глава книги, имеет существенно более сложный характер. Число посылок произвольно, да и сами посылки сложнее — ими служат закономерности типа запрета, которые могут связывать более чем два признака.

К счастью, сейчас мы обладаем мощным аппаратом теории булевых функций, развитой в первую очередь для решения задач логического проектирования разнообраз-

разных управляющих и вычислительных устройств, к которым относится и известная всем электронная вычислительная машина. Применение этого аппарата к логическому выводу позволяет успешно решать две задачи: преобразовывать систему закономерностей к более удобному виду, например, минимизируя число закономерностей, и находить строго формально все логические следствия данной системы.

Это — комбинаторные задачи, и их решение связано с ветвлением решающих процессов, с перебором вариантов, число которых быстро растет при усложнении системы закономерностей. Такой перебор неизбежен, но его можно сокращать до приемлемой величины, позволяющей решать многие практические задачи распознавания, часто прибегая к помощи электронных вычислительных машин.

Обо всем этом речь пойдет ниже.



СВЯЗИ В ПРОСТРАНСТВЕ ПРИЗНАКОВ

Как мы уже отмечали, понятие закономерности в общем случае весьма неопределенно. Однако его можно уточнять, рассматривая частные способы описания изучаемых объектов и образуемых ими классов.

Допустим, что интересующие нас объекты описываются в некоторой системе признаков, которые могут быть присущи либо нет конкретным объектам. Такие описания широко используются в различных справочниках, определителях, классификаторах.

Например, при описании промысловых рыб СССР* рассматриваются «преимущественно внешние признаки, бросающиеся в глаза и позволяющие легко отличить данный вид от близких. Таковы, например, выступающие рыло и нижний рот у рыбца; вытянутое мечевидное рыло у севрюги; зазубренные лучи в плавниках у карпа и карася...» Сельдевая акула описывается так: «Тело торпедообразное, покрыто плакоидной чешуей (шипиками). Спинных плавников два: первый — большой, впереди брюшных, второй — маленький, расположен над анальным. Колючек у спинных

* Промысловые рыбы СССР. М., 1949.

плавников нет. Жаберные отверстия велики, числом 5. Брызгальца малы или их нет вовсе. Мигательной перепонки нет. Рото-носовых бороздок нет. Скелет хрящевой».

Представим себе множество всех признаков, фигурирующих в описаниях рыб самых различных родов и видов. Чтобы определить, к какому из них относится пойманная рыба, можно предварительно составить ее описание по всем этим признакам, заполнив своеобразную анкету, где в каждой графе проставляется ответ на вопрос, обладает ли рыба соответствующим признаком — Да или Нет. Сопоставляя эту анкету с описаниями родов и видов, нетрудно определить таксономическую принадлежность рыбы.

Иногда подобное распознавание имеет чрезвычайно важное значение и в повседневной жизни. Однако не всегда удастся получить точный ответ на поставленный вопрос. Скажем, когда требуется определить, съедобен ли найденный гриб. Помимо Да и Нет возможен и третий вариант: Неопределенность. В данном случае неопределенность рекомендуется считать запретом (простое правило, нарушение которого может привести к несчастным случаям, подчас со смертельным исходом). В некоторых случаях имеет смысл сохранить понятие неопределенности. Вполне реален еще один вариант, когда четко отмечается отсутствие признака, а его наличие — неопределенно (Достоверное Нет и Вероятное Да).

Выбор каждого из этих трех вариантов (Да и Нет; Да, Нет и Неопределенность; Достоверное Нет и Неопределенное Да) зависит не только от характера объектов, но и от ряда субъективных факторов. Скажем, от «цены» за ошибку. В случае с определениями съедобности грибов цена за ошибку может быть слишком велика — отравление — при весьма скромном возможном выигрыше — полакомиться грибами. А когда речь идет о поисках полезных ископаемых, то ситуация иная: цена выигрыша велика, а неудача влечет за собой сравнительно небольшие экономические потери. В этих случаях имеет смысл различать Достоверное Нет и Неопределенное Да.

Так или иначе, распознавая некоторый объект, полезно составить его описание в пространстве некоторых признаков, получить проекцию объекта на эти признаки, его тень. Простота такого способа соответствует его

универсальности. Действительно, в терминах признаков можно описывать и болезни человека, изучаемые медициной, и химические соединения, и неисправности технических устройств, и залежи полезных ископаемых, и многое другое. Следует, однако, помнить: «беспорядочное выдвижение различных признаков не позволит определить различия предметов», как утверждали древние философы *. Выбор подходящей системы признаков для описания некоторого класса объектов — далеко не простая задача. Так, анкета, заполненная кандидатом на замещение вакантной должности научного сотрудника в некотором научно-исследовательском институте, дает слабое представление о деловых качествах кандидата.

При составлении эффективной системы признаков следует избегать двух крайностей: избыточности и недостаточности набора признаков. В первом случае важные показатели будут скрыты в массе второстепенных или малозначимых (в своеобразном информационном шуме). Во втором — не выявятся критерии для однозначного опознания конкретных объектов.

Положим, что нас интересует некоторый класс объектов, описываемых в заранее выбранной системе признаков, хорошо обоснованной, — описание в этой системе достаточно полно отражает любой конкретный объект. Условно отождествим объект с его описанием, заменяя объект присущей ему комбинацией некоторых признаков. Если признаков выбрано достаточно много, то при такой замене объекты не теряют своей индивидуальности.

Рассматриваемое множество объектов обозначим через U , полагая, что оно состоит из отдельных элементов, обозначаемых через u_i : $U = \{u_1, u_2, \dots, u_m\}$. Множество всех признаков, используемых при описании этих объектов, обозначим через $S = \{s_1, s_2, \dots, s_n\}$. Множество всех объектов, обладающих некоторым конкретным признаком s_j , обозначим через U_j , называя его классом с признаком s_j , а его дополнение, т. е. множество всех объектов, не обладающих признаком s_j , — через $\bar{U}_j$.

Например, U может представлять множество всех минералов, $s_1, s_2, s_3, \dots$ — такие признаки, как блеск, прозрачность, содержание серы в минерале и т. д.; U_3 —

* Древнекитайская философия. М., 1973. Т. 2.

множество всех серосодержащих минералов, а U_3 — множество минералов, не содержащих серы.

При исследовании реальных объектов мощность множества U , т. е. число элементов в нем, оказывается обычно очень большой, и, как правило, известна лишь относительно малая его часть. Предположим, однако (чисто умозрительно), что нам доступна вся информация об этом множестве и удалось описать каждый его элемент, перечислив признаки, которыми последний обладает, например в виде $u_1 \in (s_1, s_2, s_6)$. Это означает, что объект представляет собой комбинацию признаков s_1, s_2 и s_6 , т. е. обладает этими признаками и никакими другими. Ту же информацию можно представить иначе — строкой из нулей и единиц — булевым вектором. Символы строки соответствуют признакам $s_1, s_2, \dots$ и говорят о том, обладает ли ими объект; 1 — если обладает, 0 — если не обладает. Например, описание $u_1 = (s_1, s_2, s_6)$ можно заменить на вектор

$$1 \ 1 \ 0 \ 0 \ 0 \ 1$$

(в данном случае число признаков ограничено шестью; при большем их числе вектор дополняется справа нулями).

Пользуясь такими средствами, можно нарисовать картину, изображающую множество U в целом. Для этого следует расположить друг под другом булевы векторы, представляющие последовательно объекты u_1, u_2, u_3 и т. д., получив таким образом булеву (из нулей и единиц) матрицу. Обозначим ее через R . Матрица содержит в удобной для обозрения форме информацию об отношении принадлежности признаков объектам: если объект u_i обладает признаком s_j , то на пересечении i -й строки и j -го столбца ставится 1, в противном случае — 0.

Зачастую выбор признаков определяется возможностями их измерения. При больших ограничениях на используемые измерительные средства индивидуальность объектов может быть потеряна. Тогда отдельные строки матрицы R будут служить уже не моделями каких-то единичных объектов, а представлять их целыми группами, говоря о том, что в множестве U существуют объекты с заданными комбинациями признаков.

Представим себе такую фантастическую ситуацию. На некоторую отдаленную планету заброшен робот. За-

нимаясь различными полезными делами, он должен сам добывать себе пищу, питаясь встречающимися на планете растениями. Однако не все из них съедобны. Робот снабжен системой датчиков — измерительных устройств, с помощью которых можно определить, обладает ли конкретное растение следующими признаками:

a — электрически заряжено,

b — издает острый запах,

c — звучит,

d — ветвится,

e — светится в темноте.

Однако больше всего робота интересует признак

f — съедобно,

непосредственно неизмеряемый: только «съев» растение, робот может через некоторое время выяснить, пошло ли оно ему на пользу.

На измерение каждого признака приходится затрачивать определенные усилия. Наша задача — помочь роботу, выработав для него разумную стратегию поиска съедобных растений. Естественно, это можно сделать только на основе достоверной информации о том, какие растения произрастают на планете.

Допустим, что планета была исследована ранее и установлено, какие именно комбинации перечисленных признаков могут встречаться там у растений. Вот эти комбинации, представленные в виде матрицы *R*:

$$R = \begin{array}{c} \begin{array}{cccccc} a & b & c & d & e & f \end{array} \\ \left[\begin{array}{cccccc|c} 1 & 1 & 0 & 0 & 0 & 1 & 1 \\ 0 & 1 & 1 & 1 & 0 & 0 & 2 \\ 0 & 0 & 0 & 0 & 1 & 0 & 3 \\ 1 & 0 & 1 & 1 & 1 & 1 & 4 \\ 0 & 1 & 0 & 0 & 0 & 0 & 5 \\ 1 & 1 & 1 & 1 & 0 & 1 & 6 \\ 1 & 1 & 1 & 1 & 1 & 0 & 7 \\ 0 & 1 & 1 & 1 & 1 & 0 & 8 \\ 1 & 0 & 1 & 1 & 0 & 1 & 9 \\ 1 & 0 & 0 & 0 & 0 & 1 & 10 \\ 0 & 0 & 0 & 0 & 0 & 0 & 11 \\ 0 & 0 & 1 & 0 & 1 & 0 & 12 \\ 1 & 1 & 0 & 1 & 0 & 1 & 13 \end{array} \right] \end{array}$$

$$R = \begin{array}{c|cccccc|c} & 0 & 1 & 0 & 1 & 0 & 0 & 14 \\ & 1 & 0 & 1 & 0 & 0 & 0 & 15 \\ & 0 & 0 & 1 & 0 & 0 & 0 & 16 \\ & 0 & 1 & 1 & 0 & 1 & 0 & 17 \\ & 1 & 1 & 1 & 0 & 0 & 1 & 18 \\ & 0 & 1 & 1 & 0 & 0 & 0 & 19 \\ & 0 & 0 & 0 & 1 & 0 & 0 & 20 \\ & 1 & 0 & 0 & 1 & 1 & 1 & 21 \\ & 1 & 1 & 1 & 1 & 1 & 1 & 22 \\ & 0 & 0 & 0 & 1 & 1 & 0 & 23 \end{array}$$

Например, первая строка матрицы свидетельствует о том, что на планете встречаются растения типа u_1 , обладающие признаками a , b и f , т. е. электрически заряженные, с острым запахом и съедобные, и не обладающие при этом остальными признаками c , d и e , т. е. не звучащие, не ветвящиеся и не светящиеся. Отметим, что в число отображаемых в матрице объектов входят и объекты типа u_{11} , не обладающие ни одним из данных шести признаков. Это, конечно, не означает, что такие объекты никак нельзя распознать. Отсутствие признака — тоже признак, и в этом смысле значения 1 и 0 в матрице R равноправны. Следует, однако, помнить, что речь идет об отсутствии признаков только в данной системе координат. В нашем примере уже выявлена принадлежность объектов к растениям, т. е. заданы их общие координаты в определенном пространстве, признаков, отделяющих живое от неживого, растения от животных. Такие признаки в рассматриваемую систему не входят.

Итак, строки матрицы R — группы объектов, неразличимых в данной системе признаков. Столбцы задают классы. Так, класс A объектов с признаком a состоит из всех электрически заряженных растений; он образуется объединением групп, отмеченных значением 1 в первом столбце. Его дополнение $\bar{A}$ — множество всех незаряженных растений — получается объединением групп, отмеченных в первом столбце значением 0.

Нетрудно подсчитать, что число различных формально возможных комбинаций из шести признаков равно 64 (в общем случае 2^n для n признаков). Множество всех задающих такие комбинации булевых векторов образует

так называемое булево пространство M с довольно странными свойствами: каждая координата в этом пространстве может принимать лишь одно из двух значений 0 или 1, зато число n координат может быть большим (в нашем примере $n=6$).

Встречающиеся в растениях комбинации признаков (строки матрицы R) играют роль некоторых элементов, или точек пространства M . Множество U допустимых комбинаций является подмножеством из M : $U \subseteq M$. Назовем это подмножество областью существования объектов исследуемого класса, а его дополнение $\bar{U} = M \setminus U$ — областью запрета, поскольку данное множество образуется комбинациями признаков, которые никогда не встречаются.

Множества существования и запрета содержат одну и ту же информацию: зная одно из них, можно однозначно определить другое. Так, последовательно перебирая элементы булева пространства

$$\begin{array}{cccccc} 0 & 0 & 0 & 0 & 0 & 0, \\ 0 & 0 & 0 & 0 & 0 & 1, \\ 0 & 0 & 0 & 0 & 1 & 0, \\ 0 & 0 & 0 & 0 & 1 & 1, \\ 0 & 0 & 0 & 1 & 0 & 0, \\ \dots \end{array}$$

(для удобства элементы упорядочены по правилам традиционного двоичного кода натуральных чисел) и сравнивая их с элементами из множества U , нетрудно найти все элементы множества $\bar{U}$:

$$\begin{array}{cccccc} 0 & 0 & 0 & 0 & 0 & 1, \\ 0 & 0 & 0 & 0 & 1 & 1, \\ 0 & 0 & 0 & 1 & 0 & 1, \\ \dots \end{array}$$

Всего их будет 41 ($2^6=64$, $64-23=41$).

Описания множеств U и $\bar{U}$ в целом эквивалентны (из одного можно получить другое), однако отдельные элементы этих множеств содержат информацию существенно различного типа — с точки зрения задач распознавания. Так, знание даже одного из элементов множества запрета $\bar{U}$, т. е. информация о том, что некоторый объект не существует, позволяет иногда решать задачу распозна-

навания, в то время как аналогичные сведения о существовании некоторого объекта оказываются недостаточными для этого.

Допустим, известен один из элементов множества запрета $\bar{U}$ — элемент $0\ 1\ 1\ 0\ 1\ 1$, означающий, что нет такого растения, которое не заряжено электричеством, пахнет, звучит, не ветвится, светится и съедобно. Данную информацию можно эффективно использовать в том случае, когда роботу встретится растение, обладающее признаками b , c и e и не обладающее признаками a и d . Ясно, что это растение несъедобно, потому что в противном случае комбинация его свойства была бы задана булевым вектором $0\ 1\ 1\ 0\ 1\ 1$, а известно, что такого растения не бывает. Следовательно, встреченное растение может быть представлено только вектором $0\ 1\ 1\ 0\ 1\ 0$.

Теперь предположим, что известен один из элементов области существования U , например $0\ 1\ 1\ 0\ 1\ 0$. Можно ли решить рассматриваемую задачу, ограничившись данной информацией? Конечно, нельзя! Зная, что наблюдаемый объект обладает признаками b , c и e и не обладает признаками a и d , можно утверждать только, что он *может не обладать* признаком f . Тогда это будет известный объект $0\ 1\ 1\ 0\ 1\ 0$ — несъедобное растение. Но нельзя утверждать определенно, что наблюдаемый объект *не может обладать* признаком f , поскольку не исключается возможность существования наряду с объектом $0\ 1\ 1\ 0\ 1\ 0$ также и объекта $0\ 1\ 1\ 0\ 1\ 1$ — растение может оказаться съедобным. Разница, как видим, имеет принципиальное значение.

Поэтому формулировку закономерностей, позволяющих решать задачи распознавания, будем связывать с запретами. В нашем случае это — запрет на некоторые комбинации признаков: утверждается, что не существует объектов с такими комбинациями.

Вообще говоря, под закономерностью принято понимать некоторые связи между признаками наблюдаемых явлений, причем достаточно сильные, чтобы их можно было обнаружить при наблюдении лишь отдельных, случайно выбранных объектов из исследуемого класса при условии, что выборка будет, как говорится, представительной: объекты выбираются независимо друг от друга и в достаточном количестве.

Логичность связей в природе зачастую маскируется своеобразным «шумом» — следствием взаимодействия

большого числа признаков, не включенных в множество избранных при построении формальной модели исследуемого класса явлений, поскольку они условно относятся к малоинформативным. Тем не менее достаточно сильные закономерности пробиваются через шум.

Уточним понятие связи. Естественно считать, что признаки связаны, если не любая комбинация их значений допустима. Чем больше запретных комбинаций, тем сильнее задающая запрет связь. Самая слабая связь задается каким-либо одним элементом множества запрета $\bar{U}$ — когда именно он и запрещен. Обнаружить такую связь очень трудно. В общем случае связь охватывает несколько признаков, причем самая слабая — все признаки. Допустим, что связаны r признаков ($r \leq n$) и эта связь представляется запретом некоторой комбинации их значений. Отобразим ее уже не булевым, а троичным вектором, компоненты которого принимают значения из трехэлементного множества $\{0, 1, -\}$. При этом значением 0 или 1 будут обладать r компонент, соответствующих связываемым признакам, а остальные $n-r$ компонент получают значение «-», являющееся символом неопределенности. Связь такого вида назовем импликативной связью ранга r . Элементарная связь оказывается, таким образом, частным случаем импликативной — при $r=n$.

Чем оправдан выбор такого названия?

Обратимся к примеру. Пусть для объектов рассматриваемого класса запрещена такая комбинация значений признаков: $s_1=1, s_2=0$. Если при этом известно, что для некоторого объекта $s_1=1$, т. е. он обладает признаком s_1 , то можно с полной уверенностью утверждать, что он обладает и признаком s_2 , т. е. $s_2=1$. Короче говоря, если s_1 , то s_2 (или, что то же самое, если $s_1=1$, то и $s_2=1$). А это отношение, связывающее в данном случае признаки s_1 и s_2 , называется импликацией и записывается в виде $s_1 \rightarrow s_2$ (часто используется и такая символика: $s_1 \supset s_2$).

Импликация представляет собой формализацию связки «если ..., то ...», встречаемой в выражениях естественного языка. Однако эта связка используется в различных смыслах, что приводит к соответствующей неоднозначности толкования импликации в логике. Поэтому во избежание недоразумений уточним тот смысл, который вкладывается в рассматриваемое отношение нами.

Если признаки s_1 и s_2 находятся в отношении $s_1 \rightarrow s_2$, то это не значит, что между ними существует некоторая

причинно-следственная связь (s_1 — причина, а s_2 — следствие), хотя, вообще говоря, возможность такой связи не исключается. Это отношение означает не что иное, как невозможность одновременного наличия у некоторого предмета признака s_1 и отсутствия признака s_2 . Легко проверить, что из такой интерпретации следует, что импликация $s_1 \rightarrow s_2$ равносильна импликации $\bar{s}_2 \rightarrow \bar{s}_1$: если некоторый предмет не обладает признаком $\bar{s}_2$, то он не обладает и признаком s_1 .

Уточненную таким образом импликацию принято называть материальной в отличие от формальной импликации (секвенции). В то время как последняя применяется для выражения определенных отношений между формулами и широко используется в системах логического вывода, материальная импликация представляет фактически существующие отношения между элементами исследуемых реальных (материальных) систем, обычно выявляемые средствами, выходящими за пределы собственно логики, например опытным путем.

Заметим, что символ $\rightarrow$ применяется не только для обозначения импликации как отношения. Так же обозначается и импликация-функция, играющая роль характеристической булевой функции, принимающей значение 0 на запретной области. Точнее говоря, символ $\rightarrow$ используется в этом случае как символ операции, вычисляющей значение данной функции.

В применении к связи между признаками термин «импликативная» противопоставляется термину «функциональная». Разница в том, что при функциональной связи значения некоторых признаков, называемых аргументами, всегда определяют значение другого выбранного признака (функции), в то время как при импликативной связи — лишь иногда, при некоторых комбинациях значений исходных признаков (в рассмотренном примере при $s_1=1$). Только тогда отношение импликации «срабатывает», приводя к умозаключению определенного типа (продуктивному).

Сказанное справедливо и для случая, когда в запретной комбинации значений фигурируют более чем два признака. Только здесь левая часть выражения усложняется: она будет представлять собой элементарную конъюнкцию, обращающуюся в единицу на входящих в запретную комбинацию значениях исходных признаков. Примеры таких формул будут приведены ниже.

Естественно ожидать, что импликативные связи тем сильнее, чем меньшее число признаков они связывают. Действительно, не входящие в выражение импликативной связи ранга r признаки (их число равно $n-r$) могут образовывать 2^{n-r} различных комбинаций, которые в сочетании с единственной запретной комбинацией выделенных r признаков порождают 2^{n-r} запретных комбинаций на множестве всех признаков.

Эти запретные комбинации образуют интервал ранга r в булевом пространстве M , т. е. множество всех таких элементов пространства, которые имеют некоторые фиксированные значения r компонент и всевозможные комбинации значений остальных компонент. Именно это множество и представляется в компактной форме троичным вектором. Подставляя в последний вместо символов «—» произвольные комбинации из нулей и единиц, можно получить все элементы интервала.

Любой интервал содержит один максимальный элемент $m_{\max}$, получаемый подстановкой вместо символов «—» единиц, и один минимальный $m_{\min}$, получаемый аналогичной подстановкой нулей, а также все элементы m_i булева пространства M , оказывающиеся между ними:

$$m_{\min} \leq m_i \leq m_{\max}.$$

Неравенство носит покомпонентный характер, причем считается, что $0 < 1$. Поэтому множество рассматриваемого типа называется интервалом.

Вернемся к рассмотрению множества U допустимых комбинаций свойств у растений, растущих на планете, куда попал наш робот. Анализируя это множество, легко убедиться в существовании импликативной связи, представляемой троичным вектором

$$1 \quad - \quad - \quad 0 \quad 1 \quad -$$

и утверждающей, что не существует растений, электрически заряженных и светящихся, которые в то же время не ветвятся.

Очевидно, что эта связь эквивалентна совокупности из восьми элементарных связей, задаваемых соответственно булевыми векторами

$$\begin{array}{cccccc} 1 & 0 & 0 & 0 & 1 & 0, \\ 1 & 0 & 0 & 0 & 1 & 1, \end{array}$$

1 0 1 0 1 0,

1 0 1 0 1 1,

1 1 0 0 1 0,

1 1 0 0 1 1,

1 1 1 0 1 0,

1 1 1 0 1 1,

которые образуют интервал третьего ранга (связываются три признака) с минимальным элементом

1 0 0 0 1 0,

где все несвязанные данной связью признаки отсутствуют, и максимальным элементом

1 1 1 0 1 1,

в котором все несвязанные признаки присутствуют.

В случае, когда n велико, а r мало, имплицативную связь иногда удобнее задавать не n -компонентным трюичным вектором, а элементарной конъюнкцией — логическим произведением, где роль сомножителей играют символы связываемых признаков, причем символы признаков, не входящих в запретную комбинацию, ставятся под знак отрицания.

Так, имплицативная связь, заданная вектором $1 \text{ --- } 0 \text{ } 1 \text{ ---}$, может быть также представлена элементарной конъюнкцией $a\bar{d}e$, т. е. булевой функцией, принимающей значение 1, если $a=1$, $d=0$ и $e=1$, и принимающей значение 0 в противном случае. При этом a , d и e интерпретируются как булевы (двоичные) переменные, принимающие значение 1, если соответствующий признак входит в выбираемую комбинацию, и значение 0, если не входит.

Напомним, что выражение $a\bar{d}e$ — узаконенное сокращение формулы $a \wedge \bar{d} \wedge e$ для логического произведения показанных в ней трех сомножителей, один из которых, а именно $\bar{d}$, представляет собой отрицание переменной d , т. е. функцию, значение которой противоположно значению аргумента d . Это произведение, или конъюнкция, принимает значение 1 в том и только том случае, когда все аргументы принимают такое же значение. Другая основная функция алгебры логики — дизъюнкция, обозначаемая символом $\vee$, примет значение 1, если хотя бы один из аргументов примет это значение, и обратится в

0 лишь тогда, когда все аргументы будут иметь значение 0.

Таким образом, каждому интервалу булева пространства M соответствует своя элементарная конъюнкция, оказывающаяся характеристической функцией интервала. Она принимает значение 1 на элементах интервала и 0 за его пределами.

Как уже отмечалось, связи данного типа называются импликативными, поскольку они соответствуют импликации, т. е. отношению типа «если ..., то ...» Действительно, их всегда можно привести к такой форме.

Например, связь, заданная вектором $1 \text{ — — } 0 \text{ } 1 \text{ —}$ или соответствующей ему элементарной конъюнкцией $a\bar{d}e$, означает, что не может быть так, чтобы некоторый объект обладал признаками a и e и не обладал признаком d . Это же утверждение можно представить в форме импликации: если объект обладает признаком a и не обладает признаком d , то он не обладает и признаком e . Выразим эту импликацию более формально утверждением «если $a=1$ и $d=0$, то $e=0$ » и доведем его до компактной формулы: $a\bar{d} \rightarrow \bar{e}$.

Очевидно, что импликативная связь ранга r порождает r таких импликаций, которые равносильны между собой. Например, рассмотренная связь порождает также импликации $ae \rightarrow d$ и $\bar{d}e \rightarrow \bar{a}$. Оказываются равносильными следующие утверждения:

- 1) если растение заряжено и не ветвится, то оно не светится;
- 2) если растение заряжено и светится, то оно ветвится;
- 3) если растение не ветвится и светится, то оно не заряжено.

Такие импликации позволяют в соответствующих ситуациях однозначно определять значение некоторого признака. Но возможны и другие импликации, более общего вида, выражения $k_i \rightarrow d_i$, где k_i — некоторая элементарная конъюнкция, а d_i — элементарная дизъюнкция. Эти импликации получаются из импликативной закономерности, заданной некоторой элементарной конъюнкцией запрета, путем разделения ее на две части. Одна из них непосредственно образует элементарную конъюнкцию k_i , а другая, с инверсированными элементами, — элементарную дизъюнкцию d_i . Нетрудно подсчитать, что число таких импликаций общего вида, порождаемых

импликативной закономерностью ранга r , равно 2^r , если учесть и вырожденные случаи, в которых ранг k_i или d_i равен нулю.

В данном примере импликативная связь, заданная конъюнкцией ade , порождает восемь импликаций, включая и три приведенные выше: $ade \rightarrow 0$, $a \rightarrow d \vee \bar{e}$, $\bar{d} \rightarrow \bar{a} \vee \bar{e}$, $e \rightarrow \bar{a} \vee d$, $a\bar{d} \rightarrow \bar{e}$, $ae \rightarrow d$, $\bar{d}e \rightarrow \bar{a}$, $1 \rightarrow \bar{a} \vee d \vee \bar{e}$.

Как мы выяснили, сила импликативной связи растет с уменьшением ее ранга. Ясно, что этот ранг не может быть равен нулю, так как в таком случае все объекты оказались бы под запретом — ничего бы не существовало. Связи первого ранга также придется исключить из рассмотрения как тривиальные. Подобные связи предписывали бы всем объектам наличие или отсутствие некоторых признаков, однако в этом случае последние перестали бы служить признаками, по которым различаются хотя бы некоторые из существующих объектов.

Следовательно, среди нетривиальных наиболее сильными оказываются импликативные связи второго ранга. Каждая из них связывает некоторые два признака s_i и s_j , причем возможны четыре типа этой связи, представляемые элементарными конъюнкциями-запретами $s_i s_j$, $\bar{s}_i s_j$, $s_i \bar{s}_j$ и $\bar{s}_i \bar{s}_j$ или соответствующими им импликациями.

Такие связи исследовались еще Аристотелем в его знаменитых «Аналитиках», где излагалась строгая формальная аксиоматическая система, по-видимому, первая в истории науки, названная силлогистикой. В силлогистике рассматриваются утверждения некоторых простых типов, названные категорическими суждениями, и разработаны правила вывода одних суждений в качестве логических следствий других. Эти правила оформлены в виде силлогизмов, допускающих чисто формальное применение.

В частности, Аристотель рассматривал общеутвердительное суждение « A присуще всем B » и общеотрицательное суждение « A не присуще никакому B ». Нетрудно видеть, что эти суждения задают простые отношения между двумя классами объектов, которые можно обозначить теми же буквами A и B : первый из них составлен из всех таких объектов, которые обладают признаком A , а второй — из тех, которым присущ признак B . Первое суждение говорит о том, что класс B целиком содержится в классе A , а второе — что классы A и B не пересекаются, т. е. не содержат общих элементов. Двадцать

один век спустя Эйлер предложил удобные графические средства (рис. 2.1) для иллюстрации подобных отношений между классами (впоследствии они были названы кругами Эйлера).

Круги Эйлера делают очевидной эквивалентность общих суждений соответствующим имплицативным связям второго ранга. Действительно, суждение « A присуще всем B » равносильно утверждению «не существует

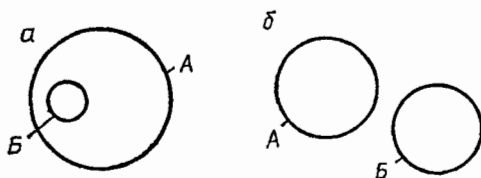


Рис. 2.1

такого объекта, который обладал бы признаком B , не обладая признаком A ». Последнее представляет собой не что иное, как связь между признаками A и B , задаваемую элементарной конъюнкцией $\bar{A}B$, налагающей запрет на комбинацию 01 значений этих признаков. Заметим, что именно в таком отношении оказываются признаки a и f исследуемых нами фантастических растений. По матрице R легко проверить, что если уж растение съедобно, то оно обязательно обладает электрическим зарядом. Аналогично может быть показано, что суждение « A не присуще никакому B » равносильно утверждению «любой объект, обладающий признаком B , не обладает признаком A », т. е. связи, задаваемой элементарной конъюнкцией AB . Что касается связи типа $\bar{A}\bar{B}$, то она легко сводится к связи $\bar{A}B$ простой перестановкой букв и их переименованием.

Связь типа $\bar{A}\bar{B}$ не рассматривается в аристотелевой силлогистике. Она означает: любой объект обладает по крайней мере одним из признаков A или B . Оказывается, ее нельзя выразить в виде простого отношения между двумя классами A и B . Сформулированное утверждение относится к множеству всех рассматриваемых объектов — так называемому универсальному классу U . Изображая эту связь с помощью кругов Эйлера (рис. 2.2), мы должны ввести и третий круг — для универсального класса U .

Заслуга Аристотеля в том, что он не ограничился исследованием и классификацией суждений (такая работа проводилась и до него), а сосредоточил основное внимание на изучении связей между суждениями. Оказалось, что последние также могут находиться в отношении импликации «если ... , то ... » Только это другого рода импликация, связывающая между собой формальные объекты, и поэтому называемая в наши дни **формальной**.

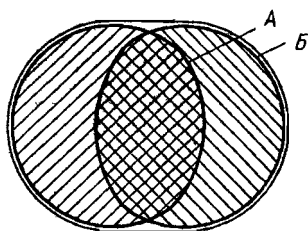


Рис. 2.2

Существование отношения формальной импликации устанавливается строго формальным путем, и первые результаты в этом направлении были получены именно Аристотелем.

Формальную импликацию называют иногда логическим следованием или логической связью. В силлогистике Аристотеля исследованы ситуации, в которых такой связью охва-

чены три суждения: из двух суждений-посылок логически вытекает третье, суждение-следствие. Например, из суждений «*A* присуще всем *B*» и «*B* присуще всем *B*» логически следует суждение «*A* присуще всем *B*». Входящие в эту тройку суждения фиксируют некоторые импликативные «материальные» связи, причем одинакового ранга — второго. Можно образовать и пары импликативных связей, вытекающих логически одна из другой, только в этом случае их ранги должны различаться. Например, из связи $s_i s_j$, охватывающей признаки s_i и s_j , очевидным образом следует связь $s_i s_j s_k$, где s_k — любой из остальных признаков, выбираемых из множества S . Будем говорить в этом случае, что связь $s_i s_j$ имплицирует связь $s_i s_j s_k$.

Импликативную связь, которая не имплицируется никакой другой импликативной связью, будем называть простой. Например, для рассматриваемого в настоящей главе множества объектов (заданного матрицей R) простой является импликативная связь $ab\bar{c}\bar{f}$. Простые связи могут играть роль ключевых, поскольку из них легко получить и все остальные импликативные связи. Это, как правило, наиболее сильные связи. Представляет практический интерес задача нахождения всех простых связей для заданного множества объектов. Ее можно ре-

шить, последовательно перебирая различные элементарные конъюнкции, образуемые символами признаков, и сравнивая их с множеством всех объектов U : если задаваемый конъюнкцией интервал не пересекается с U , то рассматриваемая конъюнкция представляет связь. Такой перебор может оказаться довольно трудоемким — ведь общее число различных элементарных конъюнкций над множеством n переменных (включая и тривиальную конъюнкцию, не содержащую ни одной буквы и обозначаемую через 1) равно 3^n . Однако его можно сильно сократить, перебирая конъюнкции в порядке возрастания их рангов. Если какая-то конъюнкция оказывается связью, рассмотрение всех конъюнкций более высокого ранга, получаемых из нее добавлением некоторых символов, излишне — все они будут представлять связи, имплицитруемые данной, и поэтому не будут простыми. С другой стороны, при таком переборе любая найденная связь может имплицитроваться только обнаруженными ранее, что существенно облегчает проверку ее простоты.

Посмотрим, как это делается все на том же примере, когда множество всех объектов U задается матрицей

$$R = \begin{array}{c|cccccc|c} & a & b & c & d & e & f & \\ \hline & 1 & 1 & 0 & 0 & 0 & 1 & 1 \\ & 0 & 1 & 1 & 1 & 0 & 0 & 2 \\ & 0 & 0 & 0 & 0 & 1 & 0 & 3 \\ & 1 & 0 & 1 & 1 & 1 & 1 & 4 \\ & 0 & 1 & 0 & 0 & 0 & 0 & 5 \\ & 1 & 1 & 1 & 1 & 0 & 1 & 6 \\ & 1 & 1 & 1 & 1 & 1 & 0 & 7 \\ & 0 & 1 & 1 & 1 & 1 & 0 & 8 \\ & 1 & 0 & 1 & 1 & 0 & 1 & 9 \\ & 1 & 0 & 0 & 0 & 0 & 1 & 10 \\ & 0 & 0 & 0 & 0 & 0 & 0 & 11 \\ & 0 & 0 & 1 & 0 & 1 & 0 & 12 \\ & 1 & 1 & 0 & 1 & 0 & 1 & 13 \\ & 0 & 1 & 0 & 1 & 0 & 0 & 14 \\ & 1 & 0 & 1 & 0 & 0 & 0 & 15 \\ & 0 & 0 & 1 & 0 & 0 & 0 & 16 \end{array}$$

$$R = \begin{array}{c|cccccc|c} & 0 & 1 & 1 & 0 & 1 & 0 & 17 \\ & 1 & 1 & 1 & 0 & 0 & 1 & 18 \\ & 0 & 1 & 1 & 0 & 0 & 0 & 19 \\ & 0 & 0 & 0 & 1 & 0 & 0 & 20 \\ & 1 & 0 & 0 & 1 & 1 & 1 & 21 \\ & 1 & 1 & 1 & 1 & 1 & 1 & 22 \\ & 0 & 0 & 0 & 1 & 1 & 0 & 23 \end{array}$$

Поскольку матрица не пуста и в каждом ее столбце содержатся как нули, так и единицы, тривиальных связей в данном случае нет. Поиск начинается с анализа элементарных конъюнкций второго ранга. При этом перебираются все 15 сочетаний из заданных шести столбцов по два ($C_6^2 = 6 \cdot 5 / 1 \cdot 2 = 15$) и для каждого из них рассматриваются четыре элементарные конъюнкции, задаваемые комбинациями значений 00, 01, 10 и 11. Связь обнаруживается в данном случае только для признаков a и f , которым соответствуют первый и последний столбцы матрицы. В этих столбцах отсутствует комбинация значений 01, т. е. оказывается пустым интервал $0 \text{ --- } 1$, не содержащий ни одного элемента из множества U . Эта связь, о которой уже говорилось выше, представляется элементарной конъюнкцией $\bar{a}f$: если объект обладает признаком f , то он обладает и признаком a ($f \rightarrow a$), или если объект не обладает признаком a , то он не обладает и признаком f ($\bar{a} \rightarrow \bar{f}$). Другими словами, все съедобные растения электрически заряжены, а все незаряженные — несъедобны. Заметим, однако, отсюда не следует, что все электрически заряженные растения съедобны.

Затем перебираются сочетания из трех столбцов (их число равно $C_6^3 = 20$) и среди них выискиваются такие, которые не содержат какой-либо комбинации значений. При этом перебор комбинаций можно сократить, не рассматривая те, которые обладают значением 0 переменной a и значением 1 переменной f , поскольку уже установлено, что они запрещены. Так находятся еще четыре связи $\bar{a}ef$, $b\bar{c}e$, $a\bar{d}e$ и $a\bar{c}\bar{f}$, поскольку в соответствующих тройках столбцов обнаруживается отсутствие комбинаций 0 1 1, 1 0 1, 1 0 1 и 1 0 0.

Аналогичным образом находятся все остальные про-

стые импликативные связи: $\bar{a}\bar{b}cd$, $\bar{b}cd\bar{f}$, $ab\bar{e}\bar{f}$, $a\bar{b}d\bar{f}$, $\bar{b}cd\bar{f}$, $ab\bar{d}\bar{f}$, $\bar{b}c\bar{d}e$, $a\bar{b}e\bar{f}$, $ad\bar{e}\bar{f}$, $\bar{b}ce\bar{f}$, $a\bar{b}\bar{c}d\bar{e}$ и $\bar{b}\bar{c}d\bar{e}\bar{f}$.

Совокупность полученных элементарных конъюнкций описывает всю область запрета $\bar{U}$: любой ее элемент принадлежит хотя бы одному из интервалов, задаваемых этими конъюнкциями. Поэтому, связав их оператором дизъюнкции, т. е. построив логическую сумму рассматриваемых конъюнкций, получим формулу для характеристической функции φ множества запрета $\bar{U}$ — такой булевой функции, которая принимает значение 1 на множестве $\bar{U}$ и значение 0 на дополнительном множестве U . Такая формула известна в булевой алгебре как дизъюнктивная нормальная форма (ДНФ). Каждый ее член представляет собой простую импликанту — элементарную конъюнкцию со следующими двумя свойствами: во-первых, она имплицирует функцию φ (если конъюнкция принимает значение 1, то и функция тоже), во-вторых, она не имплицирует никакой другой конъюнкции, обладающей первым свойством. Если в ДНФ входят все простые импликанты (как в данном случае), она называется сокращенной.

Заметим, что, говоря об импликации сейчас, применительно к некоторым функциям или представляющим их формулам, мы снова имеем в виду не материальную импликацию, использованную ранее для представления определенного типа связей между признаками, а формальную, т. е. отношение между формулами. Такое отношение обычно задается выражением $\varphi_1 \Rightarrow \varphi_2$ со следующим смыслом: если функция (или формула) φ_1 примет значение 1, то и φ_2 примет это же значение, что справедливо для любой комбинации значений переменных, входящих в данные формулы. Естественно, что при $\varphi_1 = 0$ ни одно из значений формулы φ_2 не запрещается.

Полученная сокращенная ДНФ, описывающая область запрета, имеет для рассматриваемого примера следующий вид

$$\begin{aligned}\varphi = & \bar{a}\bar{f} \vee \bar{d}e\bar{f} \vee \bar{b}c\bar{e} \vee \bar{a}d\bar{e} \vee \bar{a}\bar{c}\bar{f} \vee \bar{a}\bar{b}cd \vee \bar{b}cd\bar{f} \vee ab\bar{e}\bar{f} \vee \\ & \vee a\bar{b}d\bar{f} \vee \bar{b}c\bar{d}\bar{f} \vee ab\bar{d}\bar{f} \vee \bar{b}c\bar{d}e \vee a\bar{b}e\bar{f} \vee ad\bar{e}\bar{f} \vee \\ & \vee \bar{b}ce\bar{f} \vee a\bar{b}\bar{c}d\bar{e} \vee \bar{b}\bar{c}d\bar{e}\bar{f}.\end{aligned}$$

В качестве более наглядного способа описания области запрета можно рекомендовать матричный. При

этом используется троичная матрица (обозначим ее через Φ), каждая строка которой задает соответствующую простую импликативную связь, т. е. простую импликанту функции φ , в виде троичного вектора:

$$\Phi = \begin{array}{cccccc|c} & a & b & c & d & e & f & \\ \hline - & 0 & - & - & - & - & 1 & 1 \\ - & - & - & - & 0 & 1 & 1 & 2 \\ - & - & 1 & 0 & - & 1 & - & 3 \\ 1 & - & - & - & 0 & 1 & - & 4 \\ 1 & - & - & 0 & - & - & 0 & 5 \\ 0 & 0 & 1 & 1 & - & - & - & 6 \\ - & 0 & 1 & 1 & - & 0 & - & 7 \\ 1 & 1 & - & - & 0 & 0 & - & 8 \\ 1 & 0 & - & 1 & - & 0 & - & 9 \\ - & 0 & 1 & 0 & - & 1 & - & 10 \\ 1 & 1 & - & 0 & - & 0 & - & 11 \\ - & 0 & 1 & 0 & 1 & - & - & 12 \\ 1 & 0 & - & - & 1 & 0 & - & 13 \\ 1 & - & - & 1 & 0 & 0 & - & 14 \\ - & 0 & 1 & - & 1 & 0 & - & 15 \\ 1 & 0 & 0 & 1 & 0 & - & - & 16 \\ - & 0 & 0 & 1 & 0 & 1 & - & 17 \\ \hline \end{array}$$

Строки матрицы пронумерованы, чтобы, говоря о конкретных связях, можно было ссылаться на эти номера.

Импликативные связи могут быть связаны друг с другом не только попарно, но и более сложным образом. К примеру, данная матрица содержит только простые импликанты, а они по определению не могут имплицировать друг друга. Однако может оказаться, что какая-то из них имплицирует дизъюнкцию нескольких из остальных, например двух.

Исследуя такую возможность, удобно опираться на известную зависимость, связывающую отношение импликации между произвольными булевыми функциями f и g и отношение включения между соответствующими характеристическими множествами M_f и M_g — областями булева пространства M , на которых данные функции принимают значение 1. Эта связь весьма проста: если

функция f имплицирует функцию g ($f \Rightarrow g$), то множество M_f содержится в множестве M_g ($M_f \subseteq M_g$).

Например, рассмотрим строку 2, интерпретируя ее теперь как запретный интервал, соответствующий связи 2, т. е. как характеристическое множество элементарной конъюнкции $\bar{a}ef$:

— — — 0 1 1.

Легко убедиться, что этот интервал покрывается парой интервалов

0 — — — — 1,
1 — — 0 1 —,

задаваемых строками 1 и 4, т. е. содержится в их объединении. Действительно, его можно разложить на два интервала

0 — — 0 1 1,
1 — — 0 1 1,

первый из которых содержится в интервале 1, а второй — в интервале 4. Отсюда следует, что конъюнкция $\bar{a}ef$ имплицирует дизъюнкцию конъюнкций $\bar{a}f$ и $a\bar{e}$:
: $\bar{a}ef \Rightarrow \bar{a}f \vee a\bar{e}$.

Если какой-либо член сокращенной ДНФ функции Φ имплицирует дизъюнкцию нескольких из остальных, то его можно выбросить из формулы, получив при этом некоторую дизъюнктивную нормальную форму, уже не сокращенную, но задающую все ту же функцию Φ . Повторяя такую операцию, мы придем в конце концов к так называемой тупиковой, или безызбыточной, ДНФ, из которой уже ничего нельзя будет выбросить. При этом упрощении естественно каждый раз пытаться удалять очередную конъюнкцию, сравнивая ее с дизъюнкцией всех остальных. В самом деле, если она имплицирует дизъюнкцию лишь некоторых из них, то она должна имплицировать и дизъюнкцию всех их. С другой стороны, если анализируемая конъюнкция не имплицирует дизъюнкцию всех остальных, то она не имплицирует и никакую дизъюнкцию, образуемую некоторыми из них.

Переходя от простых импликант к соответствующим интервалам, можно утверждать, что какой-либо из интервалов, представленных в матрице Φ , можно выбросить, если он содержится в объединении всех остальных интервалов. При этом есть смысл рассматривать лишь

такие интервалы, которые пересекаются с данным, и не рассматривать такие, которые не имеют общих элементов с данным — в последнем случае троичные векторы, задающие сравниваемые интервалы, будут ортогональны, т. е. в некоторой компоненте один из них будет обладать значением 0, а другой — значением 1.

В соответствии со сказанным, задаваемый строкой 1 интервал сравнивается лишь с неортогональными строками матрицы Φ — ее минором

	<i>a</i>	<i>b</i>	<i>c</i>	<i>d</i>	<i>e</i>	<i>f</i>	
—	—	—	—	0	1	1	2
—	1	0	—	—	1	—	3
0	0	1	1	—	—	—	6
—	0	1	0	—	—	1	10
—	0	1	0	1	—	—	12
—	0	0	1	0	1	—	17

(минором матрицы, как известно, называется ее часть, образованная пересечением некоторых строк с некоторыми столбцами).

Объединение отображенных в этом миноре интервалов содержит рассматриваемый, если любой элемент последнего принадлежит в то же время некоторому из объединяемых интервалов. Перебор различных элементов указанного интервала сводится к подстановке в задающий его троичный вектор разнообразных комбинаций значений 0 и 1 вместо значений «—». Проверку отношения в целом можно представить как поиск среди получаемых в результате этой подстановки булевых векторов такого, который был бы ортогонален всем строкам минора. Поиск можно несколько упростить, не анализируя столбцы, в которых исходный вектор обладает значением 0 или 1 (по этим компонентам сравниваемые векторы не могут быть ортогональны). Таким образом, проверка сводится к поиску булева вектора, ортогонального всем строкам минора, получаемого удалением упомянутых столбцов: если такой вектор не существует, то рассматриваемая элементарная конъюнкция имплицитно дизъюнкцию стальных конъюнкций, если же существует — не имплицитно.

Подводя итоги, сформулируем необходимое и доста-

точное условие возможности удаления некоторой элементарной конъюнкции из ДНФ: не существует булев вектора, ортогонального любой строке минора, полученного из задающей эту ДНФ троичной матрицы путем удаления строки, соответствующей анализируемой элементарной конъюнкции, и всех ортогональных ей строк, а также столбцов, в которых данная строка имеет значение 0 и 1.

В данном случае — при анализе строки 1 — исследуется минор

<i>b</i>	<i>c</i>	<i>d</i>	<i>e</i>	
—	—	0	1	2
1	0	—	1	3
0	1	1	—	6
0	1	0	—	10
0	1	0	1	12
0	0	1	0	17

для которого легко найти ортогональный булев вектор. Например, его можно «сконструировать», используя в качестве первой компоненты значение 1 (в этом случае он будет ортогонален четырем последним строкам), а в качестве последней — значение 0 (тогда он будет ортогонален двум первым строкам). Остальные компоненты могут быть произвольными.

Это означает, что строку 1 из матрицы Φ выбросить нельзя — она обязательна.

А вот вторую строку удалить можно, поскольку для соответствующего ей минора

<i>a</i>	<i>b</i>	<i>c</i>	
0	—	—	1
—	1	0	3
1	—	—	4
—	0	1	10
—	0	1	12

(1)

нельзя найти ортогонального вектора.

Так, последовательно перебирая строки матрицы Φ и выбрасывая их, если это возможно, придем в конце концов к остатку

a	b	c	d	e	f	
0	—	—	—	—	1	1
—	1	0	—	1	—	3
1	—	—	0	1	—	4
1	—	0	—	—	0	5
—	0	1	1	—	0	7
1	1	—	—	0	0	8
—	0	1	0	—	1	10
—	0	0	1	0	1	17

(2)

представляющему безызбыточную ДНФ

$$\Phi = \bar{a}\bar{f} \vee \bar{b}\bar{e} \vee \bar{a}\bar{d}\bar{e} \vee \bar{a}\bar{c}\bar{f} \vee \bar{a}\bar{c}\bar{d}\bar{f} \vee \bar{a}\bar{b}\bar{e}\bar{f} \vee \bar{b}\bar{c}\bar{d}\bar{f} \vee \bar{b}\bar{c}\bar{d}\bar{e}\bar{f}.$$

Эта ДНФ описывает все ту же область запрета $\bar{U}$, но в заметно более компактной форме.

Возникает вопрос: достигается ли таким путем наиболее компактное описание области запрета $\bar{U}$?

Очевидно, что выбор дизъюнктивных нормальных форм в качестве средств описания уже уменьшает в какой-то мере возможности минимизации: могут встретиться ситуации, когда более компактными окажутся описания какого-то другого типа. Но даже если мы и ограничим средства описания, стараясь найти простейшую среди различных ДНФ, представляющих заданную область $\bar{U}$, окажется, что это не так-то просто сделать. Достижимый при упрощении матрицы Φ результат зависит от того, в каком порядке просматриваются ее строки. При разных способах образования «очереди» строк могут получаться различные тупиковые, или безызбыточные, ДНФ, содержащие иногда различное число членов. Задача нахождения среди этих ДНФ наипростейшей оказывается довольно сложной. Она исследуется в теории булевых функций.

Однако, как правило, получаемая описанным выше способом безызбыточная ДНФ является хорошим приближением к минимальной ДНФ. Будем называть ее ДНФ-описанием области запрета $\bar{U}$. В ней непосредственно перечисляются лишь некоторые из простых импликант, задающих соответствующие импликативные связи. Но из нее легко получить и все остальные простые импликанты с помощью операции «обобщенного склеива-

ния», предложенной русским логиком П. С. Порецким на рубеже XIX и XX веков (правда, в несколько иной форме).

Эта операция применима к так называемым смежным строкам, т. е. таким, которые обладают противоположными значениями (0 и 1) ровно в одной компоненте, как, например, строки 1 и 4 в матрице (1). Заключается операция в том, что образуется новая строка, в которой упомянутая компонента получает значение «—», а остальные строки значение «—», если они обе имеют такое же значение в сравниваемых строках, или значение 0 или 1, если хотя бы одна из строк имеет в соответствующей компоненте это значение.

Так в результате обобщенного склеивания строк 1 и 4 восстанавливается выброшенная ранее строка 2 со значением — — — 0 1 1.

Знание импликативных связей позволяет иногда предсказывать наличие либо отсутствие некоторого признака у объекта, о котором известно, что он обладает (или не обладает) некоторыми другими признаками. Если, скажем, известно, что объект из множества U , рассматриваемого в нашем примере, обладает признаком b , но не обладает признаком c (другими словами, $b=1$, $c=0$), то, опираясь на связь, представленную строкой 3, можно утверждать, что этот объект не обладает признаком e ($e=0$). Однако для других значений переменных b и c , например при $b=1$ и $c=1$, значение e остается неопределенным — возможно, что $e=1$, но может быть, что $e=0$.

Более определенными в этом смысле являются функциональные связи, явное предпочтение которым оказывается в областях естествознания, пользующихся классическим математическим аппаратом функционального анализа, прежде всего в физике.

Связь функциональна, если по значениям одних величин, называемых аргументами, однозначно определяются значения других — функций. В нашем случае — когда множество U всех объектов задано булевой матрицей R — это означает, что равным по значению строкам минора матрицы R , образованного столбцами, соответствующими аргументам, должны отвечать также равные значения в столбце, соответствующем признаку-функции. В противном случае предполагаемая функциональная связь не существует.

Если некоторый признак не является функцией всех остальных признаков, то он не может быть функцией и никакого другого подмножества из остальных признаков. А он не является функцией всех остальных, если в матрице R существует пара строк, различных по значению только в той компоненте, которая соответствует рассматриваемому признаку. Отсюда следует довольно простой способ выяснения, существует ли вообще какая-либо функциональная связь между признаками. Способ основан на попарном сравнении строк матрицы R и выявлении среди них соседних по значению, т. е. различающихся значением ровно в одной компоненте. Соответствующий такой компоненте признак не может быть функцией каких-либо других признаков.

Так, в рассматриваемом примере строки 1 и 10 оказываются соседними по компоненте b , и это значит, что признак b не является функцией остальных. Действительно, зная, что некоторый объект обладает признаками a и f и не обладает признаками c , d и e , нельзя однозначно ответить на вопрос, обладает ли этот объект признаком b . Аналогичные выводы можно сделать и для остальных признаков: например, для признака a , сравнивая строки 7 и 8, для признака c — строки 11 и 16, для признака d — строки 11 и 20, для признака e — строки 3 и 11, для признака f — строки 7 и 22. Так устанавливается полное отсутствие функциональных связей в данном примере.

Отсюда следует, что знание значений всех признаков, кроме одного, не гарантирует возможность вычисления последнего. Например, если робот встретит растение, которое и электрически заряжено, и пахнет, и звучит, и ветвится, и светится в темноте, то он не сможет определить, съедобно ли оно, пока не попробует.

Вообще говоря, каждая пара объектов u_i и u_j из множества U уже накладывает определенные ограничения на возможные функциональные связи между признаками. Если этим объектам присущи признаки, образующие соответственно множества S_i и S_j , то можно утверждать, что признаки из множества $S_i \oplus S_j$ не могут быть функциями остальных признаков, не принадлежащих этому множеству. Здесь через $S_i \oplus S_j$ обозначена сумма множеств S_i и S_j , т. е. множество, составленное из таких элементов, которые принадлежат или множеству S_i , или множеству S_j , но не обоим вместе.

Например, сравнивая строки 1 и 2 матрицы

<i>a</i>	<i>b</i>	<i>c</i>	<i>d</i>	<i>e</i>	<i>f</i>
1	1	0	0	0	1,
0	1	1	1	0	0,

задающие соответственно множества $\{a, b, f\}$ и $\{b, c, d\}$, легко показать, что сумма этих множеств $\{a, c, d, f\}$ — ее можно представить строкой

$$1 \quad 0 \quad 1 \quad 1 \quad 0 \quad 1$$

— содержит такие признаки, которые не могут быть функциями признаков *b* и *e*.

Рассмотрим два способа выявления функциональных связей между признаками на заданном множестве объектов *U*. Первый из них прямой: последовательно перебираются все признаки и для каждого из них выясняется, не может ли он оказаться функцией других, а если может, то каких именно. При этом отыскиваются минимальные конфигурации признаков-аргументов. Второй способ основан на попарном сравнении объектов и учете извлекаемых из этих пар ограничений на возможные функциональные связи.

Пусть, например, множество объектов *U* задано следующей матрицей:

$$R = \begin{array}{c} \begin{array}{cccccc} a & b & c & d & e & f \end{array} \\ \left[\begin{array}{cccccc} \overline{1} & 0 & 0 & 1 & 1 & \overline{0} \\ 0 & 1 & 0 & 1 & 1 & 1 \\ 0 & 1 & 1 & 0 & 0 & 1 \\ \underline{1} & 1 & 1 & 0 & 0 & \underline{0} \end{array} \right] \begin{array}{l} 1 \\ 2 \\ 3 \\ 4 \end{array} \end{array}$$

Начнем с признака *a*, соответствующего первому столбцу матрицы. Сравнивая этот столбец с остальными, видим, что признак *a* не является функцией признаков *b*, *c*, *d*, *e*, поскольку в строках 3 и 4 первый столбец обладает различными значениями, а соответствующие им комбинации значений в четырех последующих столбцах одинаковы. Однако различным значениям в первом столбце соответствуют различные значения в последнем, следовательно, признак *a* представляет собой некоторую функцию признака *f*, в данном случае — инверсию.

Можно представить признак *a* и как функцию любого

набора аргументов, включающего f , но все аргументы, кроме f , окажутся при этом фиктивными. В таком случае говорят, что зависимость от них несущественна.

В свою очередь f есть функция от a , и каждый из признаков c , d и e является некоторой функцией от любого из этих же признаков. Более сложной зависимостью с другими признаками связан признак b . Он может быть представлен как функция на следующих наборах аргументов: $\{a, c\}$, $\{a, d\}$, $\{a, e\}$, $\{c, f\}$, $\{d, f\}$ и $\{e, f\}$.

Другой способ основан, как уже говорилось, на попарном сравнении строк матрицы R . Разумеется, он приводит к тем же результатам, но иногда быстрее. Например, при сравнении строк 3 и 4 (объектов u_3 и u_4) образуется строка

$$1 \quad 0 \quad 0 \quad 0 \quad 0 \quad 1,$$

представляющая множество признаков $S_3 \oplus S_4 = \{a, f\}$. Из этой строки непосредственно следует, что ни a , ни f не могут быть функциями признаков b , c , d и e . Удобно дифференцировать данный запрет, задав его парой тричных строк

$$\begin{array}{cccccc} 1 & 0 & 0 & 0 & 0 & - \\ - & 0 & 0 & 0 & 0 & 1 \end{array},$$

— отдельно для признаков a и f .

Аналогичное рассмотрение всех пар строк матрицы R приводит к следующей булевой матрице:

	a	b	c	d	e	f
1	1	0	0	0	0	1
1	1	1	1	1	1	1
0	0	1	1	1	1	0
0	1	1	1	1	1	0
1	0	1	1	1	1	1
1	0	0	0	0	0	1

Заметим, что вторая строка тривиальна — в ней говорится, что рассматриваемые переменные не константы, т. е. для каждого признака найдется объект с этим признаком и объект без него.

Дифференциация выявленных запретов приводит к тричной матрице

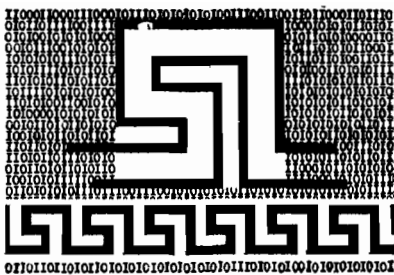
$$Q = \begin{array}{c|cccccc|c} & a & b & c & d & e & f & \\ \hline 1 & 1 & 0 & 0 & 0 & 0 & - & 1 \\ 2 & - & 1 & 0 & 0 & 0 & - & 2 \\ 3 & 0 & 1 & - & - & - & 0 & 3 \\ 4 & 0 & 0 & 1 & - & - & 0 & 4 \\ 5 & 0 & 0 & - & 1 & - & 0 & 5 \\ 6 & 0 & 0 & - & - & 1 & 0 & 6 \\ 7 & - & 0 & 0 & 0 & 0 & 1 & 7 \\ \hline \end{array}$$

Из матрицы Q можно извлечь всю информацию о функциональных связях между признаками на заданном множестве объектов и представить ее в более удобном виде. Для этого достаточно для каждого признака найти все минимальные наборы, составленные из других признаков и не содержащиеся в каком-либо из запретных для данного признака наборов, показанных в матрице Q . Полученные таким образом наборы как раз и будут наборами нефиктивных аргументов для рассматриваемого признака-функции.

Так, из строки 1 матрицы Q вытекает, что признак a есть функция признака f , а из строк 2 и 3 следует, что признак b может быть определен как функция любых двух аргументов, один из которых выбирается из множества $\{a, f\}$, а другой — из множества $\{c, d, e\}$.

Сравнивая функциональные связи с импликативными, нетрудно заметить, что они представляют собой более частный тип связей. Например, в рассмотренном ранее случае оказалось, что каких-либо функциональных связей между признаками не существует, хотя импликативных связей было довольно много.

Любая функциональная связь может быть выражена некоторой комбинацией импликативных. Так, в последнем примере функциональная связь между признаками a и f сводится к паре импликативных, задаваемых элементарными конъюнкциями af и $\bar{a}\bar{f}$, представляющими запретные наборы. В общем случае любая функциональная связь между признаком s_i и признаками $s_j, s_k, \dots, s_l$ (допустим, что число последних равно t) разлагается на 2^t импликативные связи, задаваемые трюичными векторами-запретами. В них комбинации значений компонент с номерами $j, k, \dots, l$ попарно различны, а компонента с номером i принимает произвольное значение — обратное значению функции.



ЗАКОНОМЕРНОСТИ И ИХ ВЫЯВЛЕНИЕ

Нарисованная в предыдущей главе картина идеализирована. Мы предположили, что реальные объекты внешнего мира могут быть описаны посредством признаков $s_1, s_2, \dots, s_n$. Заменили объекты их логическими описаниями — булевыми вектор-константами длиной n , а затем представили совокупность всех реальных объектов как некоторое множество таких векторов. Однако и эта модель недоступна для непосредственного исследования при решении многих практически важных задач ввиду невозможности ее полного построения. Обычно удается составить лишь некоторый фрагмент этой модели. Проводя исследования, мы чаще всего наблюдаем лишь ограниченное число реальных объектов. Нередко и этого нельзя сделать по каким-либо причинам, и тогда объекты наблюдаются частично, а некоторые компоненты представляющих их векторов получают неопределенные значения — остается неизвестным, обладает ли объект соответствующим признаком.

Такая ситуация типична, например, для медицинской диагностики. В отличие от машины человека нельзя без катастрофических последствий разложить на составные части, проведя с ними полный комплекс

исследований, желательных для достоверного определения дефектов. Более формализованный пример — из квантовой механики: принципиальная невозможность одновременного точного измерения координаты и импульса элементарной частицы (чем более тщательно измеряется один из этих параметров, тем сильнее расплывается значение другого, что отражено в соотношении неопределенности Гейзенберга).

Допустим, что объекты наблюдаются полностью (в пространстве признаков M), но число рассмотренных таким образом объектов, для которых строятся векторные модели, существенно меньше общего числа реальных объектов. Другими словами, предположим, что в результате выполненного эксперимента обнаружена и точно охарактеризована лишь некоторая часть множества U , представляющего в целом все существующие объекты изучаемого класса. Назовем эту часть множеством обнаруженных объектов и обозначим через E . Оно намного меньше множества U : $|E| \ll |U|$. Так, изучая популяцию комаров в некотором заболоченном районе, энтомолог не может измерить параметры каждого комара — он ловит и подвергает детальному исследованию лишь ничтожно малую их часть. Геолог-разведчик отбирает для последующего лабораторного анализа ограниченное число образцов минералов из интересующей его местности или даже из конкретного разреза. Такова обычная методика естественно-научных исследований.

Зная множество U в целом, мы смогли бы решать многие задачи, представляющие несомненный практический интерес. К их числу относятся и задачи распознавания, т. е. уточнения свойств новых объектов по частичным наблюдениям. Как было показано, такие задачи могут успешно решаться на основе знания связей между признаками. А эти связи уже заданы в самой структуре множества U , нужно лишь представить их в явной форме.

Но можно ли решать те же самые задачи, зная лишь множество E , обладая информацией об относительно малой части множества U ? Вопрос принципиально важный. Именно с такой ситуацией и приходится сталкиваться при исследовании большинства практических проблем.

Оказывается, решение возможно, хотя и с определенными оговорками. И в данном случае можно опираться на знания о некоторых связях между признака-

ми, хотя эти знания неизбежно будут иметь характер гипотез.

Что отражают связи между признаками? Они эквивалентны утверждениям о несуществовании объектов с некоторыми комбинациями свойств. А такие утверждения можно делать с полной определенностью, лишь обладая полной информацией о множестве U . Если наблюдению доступна лишь часть существующих объектов, представляемая множеством E , причем $E \subset U$ и к тому же $|E| \ll |U|$, никогда нельзя сказать с полной уверенностью, что объект с некоторой конкретной комбинацией признаков не существует. Остается только выдвигать более или менее правдоподобные предположения о несуществовании таких объектов.

В конце XVIII века в мире зоологов произошла сенсация: в Австралии был открыт утконос — покрытое шерстью млекопитающее с «утиным клявом» и лапами с перепонками. До этого сочетание таких признаков в одном животном считалось недопустимым. Стали подозревать, что необычное животное несет яйца. Догадка подтвердилась почти сто лет спустя.

По сходному сценарию происходят многие открытия в естествознании. Ценность данных, подтверждающих известную закономерность, убывает по мере их накопления. Одновременно возрастает эвристическая ценность факта, с помощью которого можно опровергнуть общепринятое утверждение. Этим объясняются, в частности, настойчивые поиски вечного двигателя или фактов, опровергающих первое и второе начала термодинамики.

Из общих соображений ясно, что чем сильнее некоторая связь, чем шире соответствующая ей область запрета в пространстве признаков M , тем больше шансов, что она как-то проявится и в множестве E , и ее можно будет обнаружить, рассматривая лишь это множество.

Будем называть закономерностями именно такие, выявляемые связи.

Математику это определение может показаться недостаточно четким. Он будет прав. Легко представить себе некоторую не очень сильную связь, невыявляемую на множестве E , но такую, которая может проявиться, если продолжать наблюдения за реальными объектами, строить их логические модели и дополнять ими множество E . Однако с такой нечеткостью приходится мириться, утешая себя мыслью о том, что большинство

важнейших понятий, входящих в фундамент современной науки, определено не более строго.

Тем не менее в ходе дальнейших рассуждений мы попытаемся уточнить понятие закономерности.

Итак, будем пока считать, что, анализируя множество обнаруженных объектов E , иногда можно выявить некоторые закономерности, которые позволят затем решать задачи, касающиеся исследуемого класса объектов. Совокупность полученных таким образом закономерностей составит модель множества U , т. е. вторичную модель исследуемого класса объектов.

Обратимся вновь к нашему роботу, посланному на чужую планету. Постараемся не только научить его распознавать полезные и вредные качества встречаемых на планете предметов, но и снабдить программой самообучения.

Для простоты разобьем деятельность робота на два этапа: обучения и распознавания. На этапе обучения робот должен выявить закономерности в окружающей его действительности, на этапе распознавания — опираться на них.

Займемся пока первым из этих этапов. На этапе обучения робот, снабженный достаточно богатым набором двоичных датчиков, может наблюдать за различными объектами мира, в который он попал, и определять, какими признаками из множества $\{s_1, s_2, \dots, s_n\}$ обладает очередной наблюдаемый предмет. При этом в число признаков входят и такие, роль которых априори недостаточно ясна, и такие, важность которых с точки зрения робота несомненна, например, представляет ли объект с данным признаком опасность или может оказаться полезным.

Пусть результаты своих наблюдений робот фиксирует в форме булевой матрицы E , строки которой представляют собой логические описания отдельных наблюдаемых объектов или групп неразличимых (в данной системе признаков) объектов. Закончив свои наблюдения, робот должен как-то «осмыслить» их и сформировать на этой основе некоторое представление об окружающей его действительности: построить свою модель мира, которая поможет ему в дальнейшем вести себя разумно. Как же построить ее?

Исходя из общих соображений, высказанных ранее, можно предложить роботу такую программу обработки

экспериментальных данных, заключенных в матрице E , при которой он будет искать наиболее общие связи простейшего типа, а именно импликативные, они же элементарные запреты на комбинации значений некоторых из переменных $s_1, s_2, \dots, s_n$. Поиск должен вестись в порядке убывания силы связей: сначала будут рассмотрены группы, составленные из минимального числа переменных, а затем это число будет постепенно увеличиваться. Реализуя данную программу, робот перебирает интервалы пространства M , элементами которого служат различные комбинации значений переменных $s_1, s_2, \dots, s_n$, начиная этот перебор с интервалов первого ранга, затем переходя к рассмотрению интервалов второго ранга, потом — третьего и т. д. Среди рассматриваемых интервалов могут оказаться пустые, не пересекающиеся с множеством E . Как раз этими пустыми интервалами наиболее четко отражаются существующие в исследуемом классе объектов закономерности, т. е. достаточно сильные связи между признаками.

Действительно, если некоторые из признаков связаны между собой, то существует по крайней мере одна запретная комбинация их значений, а это означает, что в соответствующий данной комбинации интервал пространства M не попадает ни один из элементов множества U и, следовательно, ни один из элементов множества E , поскольку $E \subset U$. Интервал окажется пустым.

Обратное утверждение было бы в общем случае неверным. Если некоторый интервал оказался пустым, то это еще не значит, что задаваемая им импликативная связь действительно существует. Может быть, в нем все-таки содержатся некоторые из элементов множества U , не попавшие случайно в множество E .

В таком случае роботу придется довольствоваться выдвижением гипотезы о существовании связи, соответствующей обнаруженному интервалу, не пересекающемуся с множеством E . Гипотеза, естественно, нуждается в обосновании. Требуется оценить степень ее достоверности, или, иначе говоря, вероятность того, что гипотеза окажется ошибочной. Достоверными принято называть гипотезы, вероятность ошибочности которых настолько мала, что ею можно пренебречь. Для практического руководства приемлемы именно такие гипотезы. Они и послужат материалом, из которого робот будет строить свою модель окружающей его действительности.

Как же определить достоверность гипотезы о существовании связи, соответствующей обнаруженному пустому интервалу? Ограничимся пока рассуждениями на качественном уровне.

Вспомним, что ранг пустого интервала равен числу признаков, входящих в имплицативную связь, заданную этим интервалом. Чем меньше ранг, чем больше интервал, тем больше оснований для принятия соответствующей гипотезы. В самом деле, гипотеза окажется ошибочной лишь в том случае, когда интервал содержит некоторые элементы множества U , не попавшие случайно в множество E . Но вероятность такого события быстро падает с уменьшением ранга интервала.

Поэтому поиск связей целесообразно вести, перебирая и анализируя интервалы в порядке возрастания их ранга. Достоверность выдвигаемых при этом гипотез будет монотонно снижаться, и при некотором значении ранга интервалов поиск целесообразно прекратить, поскольку его продолжение было бы бесплодным. Как выбрать это критическое значение ранга, решим дальше.

А пока уделим внимание еще одному важному вопросу: как образуется множество E и можно ли считать, что оно достаточно хорошо представляет исследуемое множество U ? Вопрос этот отнюдь не праздный. Допустим, что на планете, куда заброшен наш робот, действуют какие-то силы, которые «подсовывают» ему некоторые элементы из множества U , выбираемые таким образом, чтобы исказить истинную картину существующих в множестве U связей. Например, роботу все время может показываться один и тот же элемент. Очевидно, что в такой ситуации деятельность робота по исследованию множества U заведомо будет обречена на неудачу.

Не следует думать, что подобная ситуация совершенно умозрительна и мыслима лишь для доверчивого робота и хитрых инопланетян. Нечто сходное происходит достаточно часто в процессе выдвижения гипотез натуралистами. Исследователи природных объектов особой сложности с индивидуальными особенностями — в науках о Земле и о жизни — имеют перед собой огромное количество, нередко миллионы разнообразных фактов. Из этого множества, вряд ли охватываемого в пределах единого цельного исследования, приходится выбирать некоторое частное, относительно небольшое подмножество, по которому и определяются те или иные законо-

мерности, обосновываются гипотезы. В редких случаях этот выбор делается случайно (желательно еще выяснить, что означает такая случайность).

Геолог, например, обычно основывается на фактах, добытых при самостоятельных исследованиях. Его личный опыт неизбежно ограничен конкретными районами и геологическими формациями. Дополнительные сведения он вынужден отбирать из литературных источников, к которым склонен относиться с доверием.

Однако каждый геологический район индивидуален. Выявленные здесь закономерности можно трактовать как общепланетные в редчайших случаях и лишь в самых общих чертах. Вдобавок исследователь обычно ориентирован — или заранее, или после начала работ — на некоторые предварительные гипотезы, которые со временем могут приобрести над ним заметную власть. Предубежденно относясь к тем или иным гипотезам — принимая или отвергая их в ходе исследования, — ученый как бы «подсовывает» себе желаемое, невольно подыскивает те факты, которые подтверждают любимую ему гипотезу, и избегает фактов, опровергающих ее. Поэтому в нем самом «действуют некоторые силы, которые подсовывают ему некоторые элементы из множества U , выбираемые таким образом, чтобы исказить истинную картину существующих в множестве U связей». Увы, сказанное о работе в данной ситуации подходит и для человека.

Еще раз повторим: все это может происходить и происходит обычно (исключая сознательный обман) невольно. Исследователь искренне стремится познать природные закономерности. Это стремление поддерживается и поощряется эмоциями. Без сильного и постоянного эмоционального напряжения вряд ли можно добиться значительных успехов в познании природы. В то же время эмоции содействуют и тем силам, которые искажают картину реальности, выдавая желаемое за действительное.

Но дело, конечно, не только в субъективном факторе. Когда сталкиваешься с неохватываемым исследованием количеством фактов (которые в свою очередь охватывают ничтожную часть реальности), то обязательно встает проблема отбора, выборки. Это, можно сказать, объективно привносимый в научное исследование субъективный фактор.

В связи с этим вспоминается высказывание Эйнштейна, повторенное Н. Винером: «Бог коварен, но не злонамерен». По мнению Н. Винера: «Это замечание — не шаблонная фраза, а заявление с глубоким смыслом, касающимся стоящих перед ученым проблем. Открытие секретов природы требует мощной и тщательно разработанной техники, однако, поскольку речь идет о неживой природе, мы можем ожидать по крайней мере одной вещи — что любой возможный наш шаг вперед не будет парирован изменением стратегии природой, намеревающейся обмануть нас и сорвать наши планы... Природа играет честно...»

Хотелось бы уточнить: тут существенно то, что следует понимать под «природой» — только лишь мир вне нас или включающий нас самих. Как будто в отличие от недавнего прошлого мы склонны считать себя неотъемлемой частью природы. В таком случае злонамеренность природы — хотя бы невольная, неосознаваемая, — пожалуй, существует, что чрезвычайно осложняет проведение объективных научных исследований. Наиболее четко это сказывается на степени достоверности гипотез и многих теорий. В. И. Вернадский предлагал принципиально отделить область научных гипотез и теорий от эмпирических обобщений, максимально приближенных к объективным научным истинам (по крайней мере на данном уровне знаний).

Однако вернемся к роботу-исследователю. Предположим, что робот на неведомой планете не склонен к самообману, а злонамеренных сил на планете не существует. Пользуясь языком теории вероятностей, положим, что множество E представляет собой случайную равновероятную выборку из множества U . Это значит, что каждый элемент множества E получается путем случайного равновероятного выбора среди элементов множества U , причем выбор каждого последующего элемента не зависит от результата выбора предыдущих.

Итак, попробуем оценить достоверность гипотезы о существовании имплицативной связи, соответствующей некоторому пустому интервалу на множестве обнаруженных объектов E .

Ее можно выразить через вероятность того, что данный интервал оказался пустым чисто случайно, в то время как на самом деле никаких связей между признаками не существует. В этом случае удобно считать,

что множество U также представляет собой случайную равновероятную выборку, в данном случае из булева пространства M — полного множества всевозможных комбинаций признаков, причем все комбинации равновероятны. Следовательно, можно считать, что множество E выбирается совершенно случайно непосредственно из булева пространства M ; при этом возможен (хотя и маловероятен) повторный выбор некоторых элементов. Такое множество будем называть случайным.

Хорошим примером случайного множества может служить совокупность номеров лотерейных билетов, на которые выпал выигрыш. Автоматический механизм этой выборки демонстрируется сейчас по телевидению; когда-то выборка осуществлялась вручную: номера представлялись на шарах, вынимаемых детьми из мешка, где шары предварительно тщательно перемешивались. Конечно, выборка переставала быть случайной, если вероятность выигрыша на некоторые наперед заданные номера искусственно повышалась. Один из способов такого мошенничества описан в романе Габриэля Гарсиа Маркеса «Осень патриарха». Шары с номерами, которые по замыслу должны быть выигрышными, заранее подогревались, а детям говорили, что если они найдут шар потеплее, то получают конфетку...

Представим множество E одноименной булевой матрицей, образованной n столбцами, соответствующими признакам $s_1, s_2, \dots, s_n$, и m строками, задающими выбранные объекты, — положим, что их число равно m . Поскольку каждый из них выбирается случайно и независимо от других из множества M , можно считать, что матрица E также оказывается случайной: каждый ее элемент принимает значение 0 или 1 с одной и той же вероятностью 0,5 и независимо друг от друга. Так как число элементов матрицы равно произведению числа строк на число столбцов, то общее число различных матриц оказывается равным 2^{mn} , причем все они равновероятны.

Согласно терминологии теории вероятностей, назовем эти различные и равновероятные исходы выбора матрицы E шансами, а вероятность интересующего нас события определим как отношение числа благоприятных шансов, при которых событие наступает, к общему числу шансов.

Выберем в пространстве M некоторый конкретный

интервал ранга k и будем рассматривать событие, заключающееся в том, что он не будет содержать элементов, общих со случайным множеством E , иначе говоря, интервал окажется пустым. Обозначим это событие через $i(m, n, k)$, где в скобках показаны параметры события: m — число элементов в случайном множестве E ; n — число признаков; k — ранг интервала. Подсчитаем вероятность $p_i(m, n, k)$ события.

Эта вероятность не зависит от того, какие именно признаки будут характеризовать интервал и какими они будут обладать значениями. Напомним, что 1 означает присутствие соответствующего признака в рассматриваемом объекте, а 0 — его отсутствие. Как признаки, так и их значения входят в нашу модель симметрично. Для удобства положим, что интервал образован нулевыми значениями первых k признаков: $s_1, s_2, \dots, s_k$. Это значит, что он окажется пустым в множестве E , если ни одна из строк матрицы E не будет содержать в первых k позициях только нули.

Подсчитаем вероятность такого события. Рассмотрим множество всех различных булевых матриц размером $m \times n$. Доля матриц, у которых первые k позиций в первой строке заполнены нулями, составляет 2^{-k} от общего числа матриц 2^{mn} . Доля остатка равна, очевидно, $1 - 2^{-k}$. Следовательно, ровно $2^{mn} \times (1 - 2^{-k})$ различных матриц имеют хотя бы одно значение 1 в первых k позициях первой строки. Точно также определяется доля тех из них, которые будут иметь хотя бы одну единицу в первых k позициях также и во второй строке. Число таких матриц оказывается равным $2^{mn} (1 - 2^{-k})^2$. Аналогичное рассмотрение последующих строк приводит к следующему результату: число матриц, содержащих не менее одной единицы в первых k позициях каждой строки, равно $2^{mn} (1 - 2^{-k})^m$.

Вероятность того, что некоторый произвольный, но конкретный интервал ранга k окажется пустым при случайно выбранном множестве E с параметрами m и n , получается при делении вычисленного таким образом числа благоприятных шансов на общее число различных матриц:

$$p_i(m, n, k) = (1 - 2^{-k})^m.$$

Допустим, при анализе множества E найден некоторый пустой интервал множества M . По приведенной

выше формуле подсчитано, что при случайном характере множества E вероятность того, что именно этот интервал окажется пустым, довольно мала, например равна 0,0001. Можно ли на основании этого утверждать, будто множество E отнюдь не случайно и в нем отражается определенная связь между соответствующими признаками? Или оснований для такого вывода нет, поскольку данный интервал оказался пустым совершенно случайно и удивляться тут нечему?

Тут мы подошли к самому ответственному этапу рассуждений.

Конечно, априорная вероятность оказаться пустым для данного конкретного интервала весьма мала. Но ведь в пространстве существует много различных интервалов. Вероятность того, что какой-то из них окажется пустым, может быть существенно выше. Более того, если не ограничивать ранг интервалов, в пространстве M обязательно найдутся пустые интервалы, поскольку множество E представляет собой лишь часть этого пространства. В конце концов можно выбрать некоторый интервал ранга n , состоящий из какого-нибудь одного элемента, не принадлежащего множеству E . Стоит ли удивляться, если какой-то интервал окажется пустым?

Однако поинтересуемся более сильными связями, ограничив сверху ранги анализируемых интервалов. Рассмотрим событие «некоторый интервал ранга k не пересекается со случайным множеством E », которое обозначим через $I(m, n, k)$. Оно представляет собой сумму всевозможных событий типа $i(m, n, k)$, касающихся конкретных интервалов: событие $I(m, n, k)$ наступает при любом событии $i(m, n, k)$, соответствующем выбору любого интервала ранга k . Очевидно, что событие $I(m, n, k)$ наступит и в том случае, когда пустым окажется какой-либо интервал меньшего ранга, так как этот интервал будет обязательно содержать несколько пустых интервалов ранга k .

Обозначим через $p_I(m, n, k)$ вероятность события $I(m, n, k)$. Ее-то и предстоит нам оценить.

При расчете вероятностей сложных событий, выражающихся каким-то образом через простые, вероятности которых известны, используются обычно две основные теоремы теории вероятностей.

Одна из них относится к сумме несовместимых, или взаимно исключающих друг друга событий (которые

никогда не могут наступить одновременно): вероятность суммы таких событий равна сумме вероятностей этих событий. Более формально утверждение выглядит так: если A и B — несовместимые события, то

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

В общем же случае, когда отношение между событиями A и B произвольно, справедлива формула

$$p(A + B) \leq p(A) + p(B).$$

В ее справедливости можно убедиться, глядя на рис. 3.1, где события A и B показаны кружками, а их вероятности — занимаемой ими частью пространства M , представленного прямоугольником. Легко сообразить, что площадь, занимаемая кружками совместно, не может быть больше суммы площадей, занимаемых кружками по отдельности.

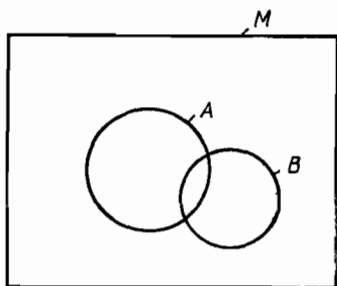


Рис. 3.1

Другая теорема применима при определении произведения независимых событий. Произведением событий A и B называется такое событие, которое наступает в том и только том случае, когда наступает и событие A , и событие B . Оно обозначается через $A \cdot B$. События считаются независимыми, если при наступлении одного из них вероятность наступления другого не изменяется. Теорема утверждает, что вероятность произведения независимых событий A и B равна произведению их вероятностей:

$$p(A \cdot B) = p(A) \cdot p(B).$$

Рассматриваемое событие может представлять собой достаточно громоздкое переплетение некоторых простых событий, вероятности которых известны, причем таких, которые находятся в сложных отношениях друг с другом. Тогда вычисление его вероятности сводится к предварительному преобразованию представляющей событие формулы таким образом, чтобы изобразить его в виде некоторой суперпозиции из произведений независимых

событий и сумм несовместимых. На этом пути приходится сталкиваться со сложными комбинаторными проблемами, решение которых не всегда удается практически довести до конца.

Именно такая ситуация возникает при рассмотрении события $I(m, n, k)$, представляющего собой сумму событий типа «интервал ранга k не пересекается с множеством E », которая берется по всевозможным интервалам такого типа. Оказывается, эти события связаны друг с другом, причем далеко не просто, поскольку взаимные пересечения интервалов образуют довольно сложную картину.

Однако нам не надо вычислять значение вероятности $p_I(m, n, k)$ на всем пространстве параметров m, n и k . Достаточно оценить эту вероятность при малых ее значениях. Только если она достаточно мала, наступление события $I(m, n, k)$ вызывает какой-то интерес и можно выдвинуть гипотезу о существовании соответствующей связи между некоторыми признаками.

В этом случае оказывается, что величина $p_I(m, n, k)$ хорошо оценивается другой величиной: математическим ожиданием $W(m, n, k)$ числа интервалов ранга k , не пересекающихся со случайным множеством E .

Как известно, математическим ожиданием случайной величины называется ее значение, усредненное по всем возможным реализациям с учетом их вероятностей — путем соответствующего взвешивания. В данном случае все реализации — различные булевы матрицы размером $m \times n$ — равновероятны, поэтому подсчет математического ожидания упрощается.

Рассмотрим, с одной стороны, множество всех 2^{mn} различных булевых матриц заданного размера, а с другой — множество всех интервалов ранга k в пространстве M . Число таких интервалов равно $C_n^k 2^k$: первый сомножитель показывает число различных сочетаний при выборе k признаков из общего их числа n , а второй — количество фиксируемых комбинаций значений для выбранных признаков. При этом

$$C_n^k = \frac{n \cdot (n-1) \cdot \dots \cdot (n-k+1)}{1 \cdot 2 \cdot \dots \cdot k}.$$

Рассмотрим далее произведение этих множеств, образуемое всевозможными парами, первый элемент кото-

рых выбирается из множества матриц, а второй — из множества интервалов. Подсчитаем число таких пар, которые находятся в отношении непересечения.

Воспользуемся подсчитанным ранее числом различных матриц, не пересекающихся с некоторым конкретным, но произвольно выбранным интервалом ранга k — оно равно $2^{mn}(1-2^{-k})^m$. Умножив данную величину на число различных интервалов ранга k , получим $2^{mn}C_n^k \cdot 2^k(1-2^{-k})^m$, а разделив этот результат на число различных матриц, найдем интересующее нас значение математического ожидания:

$$W(m, n, k) = C_n^k 2^k (1 - 2^{-k})^m.$$

Из самого способа подсчета значения $W(m, n, k)$ видно, что оно равно сумме вероятностей событий $i(m, n, k)$, взятой по всевозможным интервалам ранга k . А поскольку сумма вероятностей любых событий не может быть меньше вероятности суммы этих же событий, справедливо отношение

$$p_I(m, n, k) \leq W(m, n, k).$$

Чтобы получить более наглядное представление о величине W , вычислим (хотя бы с помощью карманного калькулятора) некоторый ряд ее значений для случая, когда число признаков равно 100, а число объектов, попадающих в множество E , — 200.

Результаты представим в табличной форме:

k	$W(200, 100, k)$	k	$W(200, 100, k)$
1	$1,24 \cdot 10^{-58}$	4	$1,56 \cdot 10^2$
2	$2,04 \cdot 10^{-21}$	5	$4,21 \cdot 10^6$
3	$3,26 \cdot 10^{-6}$	6	$3,27 \cdot 10^9$

Бросается в глаза сильная зависимость W от k , существенно упрощающая принятие решений при выборе гипотез. Когда $m=200$, а $n=100$, при ранге интервала, не превышающем 3, значение математического ожидания оказывается весьма малым, а значит (поскольку $p_I \leq W$); не бóльшим будет и значение вероятности того, что найдется интервал ранга не выше 3, не пересекающийся с множеством E , случайно выбранным из про-

странства M . Следовательно, если такой интервал все же будет обнаружен, то вполне оправдано выдвижение гипотезы о соответствующей закономерности, нашедшей отображение в этом пустом интервале. Но если $k > 3$, положение резко меняется, и нахождение пустого интервала ранга k уже не может служить основанием для гипотезы о соответствующей имплекативной закономерности, связывающей некоторые признаки.

Так, на множестве значений k находится достаточно четкая граница, отделяющая интервалы, которые целесообразно исследовать, проверяя их на пересечение с множеством E , от других интервалов, рассмотрение которых совершенно излишне, поскольку оно ни к чему не приводит. Эта граница зависит от выбора порогового значения ω величины W , а выбор в свою очередь зависит от ряда соображений, в которые не будем здесь вдаваться. В данном примере граница лежит между рангами 3 и 4. Значит, следует искать имплекативные закономерности ранга не выше 3. Исчерпывающее решение этой задачи обеспечивается перебором всех интервалов ранга 3 в булевом пространстве ста переменных. Число таких интервалов определяется по формуле $C_n^k 2^k$ и равно $C_{100}^3 \cdot 2^3 = 1\,293\,600$. Не так уж много, если исследовать интервалы с помощью вычислительной машины. Во всяком случае это существенно меньше числа всех интервалов, равного $3^n = 3^{100} \approx 5,15 \cdot 10^{47}$. Такое число не смогла бы перебрать никакая вычислительная машина!

Интересно выяснить, как зависит положение упомянутой границы от значений m и n . Некоторое представление об этой зависимости дает следующая таблица, где элементами служат максимальные значения ранга интервалов, которые стоит исследовать при соответствующих значениях m и n (данные получены для порога $\omega = 0,01$):

n	m					
	20	50	100	200	500	1000
10	1	2	3	4	5	6
30	1	2	2	3	4	5
100	1	1	2	3	4	5

Из приведенной таблицы можно сделать два вывода, справедливых по крайней мере в рассматриваемой об-

ласти параметров. Во-первых, чтобы повысить на единицу ранг отыскиваемых закономерностей, надо удвоить объем эксперимента — число объектов в множестве E . Во-вторых, примерно тот же эффект дает уменьшение в 10 раз числа признаков, в терминах которых описываются объекты.

Отметим относительный характер достоверности выдвигаемых гипотез. Пусть, к примеру, за выбираемыми объектами следят одновременно два наблюдателя. Оба они имеют дело с одной и той же выборкой из 100 объектов, но первый из них учитывает значения 30 признаков $s_1, s_2, \dots, s_{30}$, а второй сосредоточил свое внимание на первых десяти. И пусть они обнаружили некоторый пустой интервал ранга 3, например, задаваемый элементарной конъюнкцией $s_1 s_2 s_3$. Спрашивается, как могут они оценить это событие?

Руководствуясь полученной выше таблицей, второй наблюдатель может вполне обоснованно выдвинуть гипотезу о том, что конъюнкция $s_1 s_2 s_3$ представляет закономерную связь между признаками s_1, s_2 и s_3 , которые, следовательно, не могут одновременно принадлежать некоторому объекту. Но первый наблюдатель не вправе сделать такой вывод, поскольку с его точки зрения данный интервал мог оказаться пустым совершенно случайно.

Кажущийся парадокс объясняется различием моделей, в которых отражаются происходящие события. Если кто-то выиграл в лотерею легковую машину, то он с полным правом удивляется (и, конечно, радуется): с его точки зрения произошло маловероятное событие. Однако с точки зрения организаторов лотереи ничего удивительного должен же был кто-то выиграть машину.

«Везунчик» оценивает происходящее, проектируя его на свою модель, в которой априори придается особо большая значимость возможному выигрышу машины именно им, а не кем-либо другим. То, что машину могут выиграть другие игроки, его мало интересует, — все такие выигрыши сливаются для него в некоторое одно событие, вероятность которого очень мало отличается от единицы. Поэтому если все-таки выиграл он, а не они, то для него произошло действительно удивительное событие. Организаторы лотереи рассуждают иначе. Они знают, что кто-то должен выиграть, и с их точки зрения «везунчик» ничем не выделялся априори среди других обладателей

лотерейных билетов — безликих «кто-то». Поэтому они и не видят ничего удивительного в происшедшем.

Вообще говоря, каждая из 2^{mn} различных реализаций матрицы E очень маловероятна — ее вероятность равна 2^{-mn} . Однако какая-то из них все же наступает. Но, если мы не имели в виду именно эту конкретную реализацию еще до ее наступления, удивляться нечему.

Вспомним про робота, вынужденного искать себе пищу на далекой планете. Допустим теперь, что на этапе обучения робот может позволить себе не только измерять пять упомянутых ранее признаков у входящих в обучающую выборку растений, но и пробовать, съедобны ли они. Результаты экспериментов робот накапливает в виде булевой матрицы E . Когда их наберется достаточно много, отыскивает в пространстве признаков M оставшиеся пустыми интервалы ограниченного ранга и выдвигает гипотезы о соответствующих им имплекативных закономерностях. Совокупность последних как раз и составит для робота «модель мира», которой он может впоследствии пользоваться, рассматривая некоторые новые объекты и определяя свое отношение к ним. Может случиться и так, что эта модель окажется пустой — когда не удастся найти ни одного интервала с заданными свойствами. В таком случае следует признать отсутствие достаточно сильных закономерностей, которые могли бы проявиться в матрице E . Для выявления связей более высокого ранга, более «тонких», следовало бы значительно увеличить объем эксперимента, порождающего матрицу E .

Предположим, полученная роботом матрица E такова, что после «приведения подобных» среди строк она совпадет с матрицей R , рассмотренной нами ранее. Прежде матрица R представляла множество всех допустимых комбинаций признаков. Сейчас она интерпретируется как множество комбинаций, встреченных при обучении. Как было показано ранее, множество всех пустых интервалов в пространстве M в данном случае задается троичной матрицей

$$\begin{array}{cccccc} a & b & c & d & e & f \\ \begin{vmatrix} 0 & - & - & - & - & 1 \\ - & 1 & 0 & - & 1 & - \\ 1 & - & - & 0 & 1 & - \end{vmatrix} \end{array}$$

$$\begin{bmatrix} 1 & - & 0 & - & - & 0 \\ - & 0 & 1 & 1 & - & 0 \\ 1 & 1 & - & - & 0 & 0 \\ - & 0 & 1 & 0 & - & 1 \\ - & 0 & 0 & 1 & 0 & 1 \end{bmatrix}$$

Всего пустых интервалов восемь. Один из них обладает рангом 2, три — рангом 3, три — рангом 4 и один — рангом 5. Какие же из этих интервалов отражают закономерности растительного мира планеты?

Выбор определяется числом m растений, исследованных на этапе обучения, в соответствии с оценкой W и зависит также от порогового значения ω этой оценки.

Допустим, нас устраивает значение $\omega=0,001$. Тогда для того чтобы истолковать как закономерность пустой интервал ранга 2

$$0 \quad - \quad - \quad - \quad - \quad 1,$$

надо (согласно полученной выше формуле), чтобы m было не менее 39, при $m \geq 90$ закономерными можно признать также три пустых интервала ранга 3

$$\begin{array}{cccccc} - & 1 & 0 & - & 1 & - \\ 1 & - & - & 0 & 1 & - \\ 1 & - & 0 & - & - & 0, \end{array}$$

при $m \geq 192$ — еще три пустых интервала ранга 4

$$\begin{array}{cccccc} - & 0 & 1 & 1 & - & 0, \\ 1 & 1 & - & - & 0 & 0, \\ - & 0 & 1 & 0 & - & 1, \end{array}$$

наконец, при $m \geq 384$ в модель мира вносятся все пустые интервалы, включая интервал ранга 5

$$- \quad 0 \quad 0 \quad 1 \quad 0 \quad 1.$$

Если же число испытанных растений менее 39, объем эксперимента в данном случае следует признать недостаточным для выдвижения каких-либо гипотез о закономерностях.

Пусть, к примеру, робот испытал в режиме обучения сто случайно встреченных растений и на основе полученных данных обнаружил восемь перечисленных выше пустых интервалов в пространстве признаков. Тогда ему

следует ограничиться учетом лишь первых четырех из них. Моделью мира будет служить задающая их матрица

$$\begin{matrix} & a & b & c & d & e & f \\ \begin{matrix} 0 \\ - \\ 1 \\ 1 \end{matrix} & \begin{bmatrix} - & - & - & - & 1 \end{bmatrix} \end{matrix}.$$

Как же робот сможет руководствоваться в поисках пищи такой моделью? Более подробно об этом будет рассказано в следующих главах. А пока, глядя на эту конкретную матрицу, нетрудно прийти к выводу: если роботу встретится растение без признака a , то оно не будет обладать и признаком f ; если же растение обладает признаком a , но не обладает признаком c , то оно обладает и признаком f . Иначе говоря, роботу следует иметь в виду, что съедобные растения могут встретиться только среди тех, которые обладают электрическим зарядом, а если, кроме того, растение не издает звуков, то уж точно, оно съедобно.

В заключение посмотрим, как можно оценить достоверность гипотезы о функциональной закономерности.

Как и прежде, будем считать, что матрица E размером $m \times n$ заполняется нулями и единицами совершенно случайным образом, т. е. появление нуля или единицы в каждой позиции равновероятно и заполнение различных позиций происходит независимо друг от друга.

Рассмотрим событие $f(m, n, k)$, при котором некоторые столбы окажутся охваченными функциональной связью ранга k . Такое событие происходит, когда в матрице найдутся некоторые k столбцов (назовем их аргументами) и еще один столбец (функция), обладающие следующим свойством: если некоторые строки матрицы E имеют одинаковые комбинации значений в столбцах-аргументах, то они обладают одинаковыми значениями и в столбце-функции.

Если указанные столбцы будут выбраны в матрице заранее, то вероятность $p_f(m, n, k)$ того, что они окажутся в функциональной связи, можно оценить сверху величиной

$$2^{-(m-2^k)},$$

получаемой на основе следующих рассуждений.

Выберем в матрице E некоторый минор, образуемый k столбцами, претендующими на роль аргументов. Выделим в нем строки, не совпадающие по значению ни с одной из расположенных выше. Выделенные строки минора, следовательно, будут обладать различными значениями, а все остальные — совпадать по значению с какой-либо из них. Очевидно, число выделенных строк не больше чем 2^k , а остальных не меньше чем $m - 2^k$. Таким образом, строки матрицы E в целом будут разбиты на классы. В каждом из них будут находиться строки с одинаковыми комбинациями значений в столбцах-аргументах. Общее число классов не превысит 2^k .

В случае функциональной связи можно считать, что столбец-функция принимает произвольные значения в выделенных строках, если в любой другой строке его значения будут те же, что и в выделенной, попавшей в тот же класс. При интересующем нас событии все элементы рассматриваемой совокупности окажутся свободными, кроме некоторых из них, общим числом не более $m - 2^k$. А отсюда следует, что

$$p_f(m, n, k) \leq 2^{-(m-2^k)}.$$

Но нас интересует не какая-то определенная функциональная связь ранга k , а логическая сумма всех таких событий. Обозначим ее через $F(m, n, k)$. Это событие наступает тогда, когда в матрице E найдется хотя бы одна конкретная функциональная связь ранга k , независимо какая именно.

Вероятность суммы событий ограничена сверху величиной $Q(m, n, k)$ — суммой вероятностей отдельных событий. Число последних равно $C_n^k (n - k)$. Первый множитель показывает число различных выборов на роль аргументов k столбцов из матрицы E , а второй — число выборов среди остающихся столбцов одного столбца-функции. Отсюда следует, что

$$Q(m, n, k) = C_n^k (n - k) 2^{-(m-2^k)}.$$

Осталось выяснить качество оценки. Очевидно, что при $m < 2^k$ оценка оказывается слишком грубой. Но нас интересует область малых значений вероятности $p_F(m, n, k)$, так как только там можно принимать гипотезы о функциональных связях на основании проведенных наблюдений. В этом случае полученная оценка оказывается достаточно точной.

Подсчитаем значения величины Q для тех значений параметров, которые были использованы при изучении импликативных связей. Получим следующую таблицу:

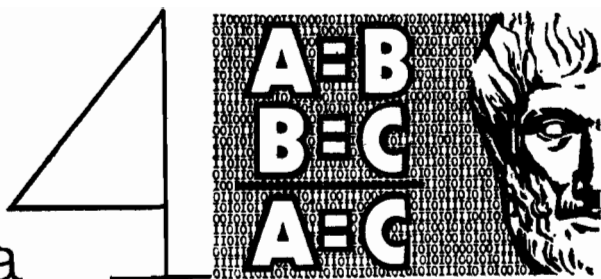
k	$Q(200, 100, k)$	k	$Q(200, 100, k)$
1	$2,46 \cdot 10^{-58}$	5	$1,91 \cdot 10^{-41}$
2	$4,83 \cdot 10^{-54}$	6	$1,29 \cdot 10^{-30}$
3	$2,50 \cdot 10^{-51}$	7	$3,16 \cdot 10^{-10}$
4	$1,54 \cdot 10^{-47}$	8	$1,23 \cdot 10^{30}$

Из анализа таблицы следует, что при $k \leq 7$ вероятность случайного образования в матрице E функциональной связи ранга k достаточно мала, чтобы в случае обнаружения такой связи можно было выдвинуть гипотезу о существовании соответствующей закономерности в том множестве, из которого выбираются строки матрицы E . Интересно заметить, что эта область значений k удовлетворяет неравенству $2^k < m$, при нарушении которого значение величины Q резко увеличивается.

Итак, при $m=200$ и $n=100$ можно вести поиск функциональных закономерностей, ранг которых не превышает семи. Определив максимальное значение ранга при других ситуациях (при пороге 0,01), получим таблицу

n	m					
	20	50	100	200	500	1000
10	2	5	6	7	8	9
30	1	4	6	7	8	9
100	—	3	5	7	8	9

Ее можно сравнить с аналогичной таблицей, полученной при рассмотрении импликативных закономерностей. При том же объеме экспериментальных данных, заключенных в матрице E , ранги функциональных закономерностей, которые есть смысл искать, заметно выше по сравнению с импликативными закономерностями. Объясняется это тем, что функциональные связи сильнее импликативных. В самом деле, функциональная связь сводится к серии импликативных, соответствующие которым интервалы покрывают упорядоченным образом ровно половину булева пространства M , превращая ее в область запрета.



СИГЛОГИСТИЧЕСКАЯ МОДЕЛЬ МИРА

Наблюдая за отдельными объектами окружающей его действительности и убеждаясь, таким образом, в их существовании, робот может строить свою модель мира, которая поможет впоследствии ему ориентироваться в этом мире.

Если объекты описываются в терминах двоичных признаков и представляются затем как булевы векторы — точки булева пространства, то наиболее простым и в то же время достаточно универсальным способом построения модели мира является поиск имплицативных закономерностей, отображаемых пустыми интервалами пространства признаков, т. е. такими интервалами, в которые не попал ни один из обнаруженных объектов.

Напомним, что при этом робот должен ограничиваться рассмотрением интервалов достаточно малого ранга. Последовательно отыскивая среди них пустые, он сокращает область возможного существования объектов, облегчая тем самым решение последующих задач распознавания. Так, из глыбы мрамора, последовательно отсекая зубилом «лишние» куски, скульптор создает скульптуру — его «модель мира».

Поскольку каждая имплицативная закономерность может быть задана в виде эле-

ментарной конъюнкции, то модель в целом можно представить как дизъюнкцию последних. Так получается дизъюнктивная нормальная форма (ДНФ), описывающая запретную область в пространстве признаков. Будем считать, что ДНФ обладает рангом, равным максимальному рангу входящих в нее конъюнкций. Каков же он может быть?

Как уже говорилось, выявляемые закономерности носят характер гипотез и могут быть приняты за достоверные лишь в том случае, когда соответствующие им обнаруженные пустые интервалы достаточно велики, так как в противном случае можно предположить, что они оказались пустыми чисто случайно. Отсюда следует, что ранги импликативных закономерностей, из которых строится модель мира, должны быть ограничены. Самые сильные закономерности должны обладать минимальным рангом.

Закономерности ранга 0 или 1 тривиальны. Первая из них запрещает существование каких бы то ни было объектов. Ясно, что она отвергается при обнаружении хотя бы одного из них. Закономерность ранга 1 представляет собой утверждение о фиктивности некоторого признака — он или присущ всем объектам, или не присущ никакому из них. Очевидно, что такой признак можно сразу же выбросить из множества признаков S , поскольку он не несет никакой информации, которую можно было бы использовать при решении задач распознавания. Другими словами, такой признак лучше не считать признаком.

Будем полагать поэтому, что в простейшем случае ранг выявляемых импликативных закономерностей ограничен сверху числом 2 (например, таким числом придется ограничиться, когда в исследуемом классе объектов, описываемых в терминах 30 признаков, обнаружено 50 объектов). В этом случае модель мира будет строиться из элементарных конъюнкций второго ранга, т. е. будет представлять собой ДНФ второго ранга.

Построение такой модели не представляет больших затруднений. Достаточно перебрать в пространстве M все интервалы второго ранга (а их число равно $C_n^2 2^2 = 2n(n-1)$) и выявить среди них пустые, не пересекающиеся с множеством обнаруженных объектов.

Импликативные закономерности второго ранга уже сопоставлялись выше с категорическими суждениями,

для которых Аристотелем и его последователями был разработан формальный аппарат, названный силлогистикой и в течение двух тысячелетий служивший фундаментом логики. Даже теперь, обратившись к словарю*, можно увидеть, что львиная доля определений посвящена именно этой замечательной формальной системе.

Займствуя определения из данного словаря, напомним читателю, что суждение — «форма мысли, в которой утверждается или отрицается что-либо относительно предметов и явлений, их свойств, связей и отношений и которая обладает свойством выражать либо истину, либо ложь». Как видно, такое определение годится для суждений в самом широком смысле этого слова. Категорические суждения, рассматриваемые в силлогистике Аристотеля, представляют собой частный, но достаточно важный случай суждений.

Категорическим называется «суждение, в котором выражается знание о принадлежности или непринадлежности признака предмету, независимо от каких-либо условий». Формально категорическое суждение представляется в виде « S есть P » или « S не есть P », где через S и P обозначаются соединяемые с помощью связки «есть» или «не есть» компоненты суждения, называемые терминами: S — предмет мысли, субъект (*subjectum*), P — то, что говорится о предмете, его свойство, или предикат (*praedicatum*). Роль связок могут играть также выражения «суть» или «не суть» — когда термин S представляется во множественной форме.

Например, суждениями являются предложения «все птицы имеют перья», «этот котенок голоден», «свинец — тяжелый металл» — в последнем из них связка «есть» заменена тире. Роль субъекта в них играют все птицы, этот котенок и свинец, а правая часть предложений — предикат — сообщает некоторые сведения об этих субъектах.

Аристотелем рассматривались четыре основных вида категорических суждений, для которых впоследствии были введены специальные обозначения:

A — общеутвердительное: все S суть P ,

E — общеотрицательное: ни одно S не есть P ,

I — частноутвердительное: некоторые S суть P ,

O — частноотрицательное: некоторые S не суть P .

* Кондаков Н. И. Логический словарь-справочник. М., 1975.

Буквы, выбранные для обозначения этих видов, легко запомнить, связав их с латинскими словами *affirmo* (утверждаю) и *nego* (отрицаю). В случае суждения общего типа для его обозначения используется первая гласная буква соответствующего слова, в случае суждения частного типа — вторая гласная.

Так, суждение «все птицы имеют перья» является общеутвердительным. Примером общеотрицательного суж-

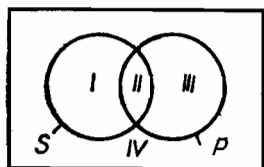


Рис. 4.1

дения может служить предложение «рожденный ползать летать не может», очевидно, эквивалентное более формальному «ни одно существо, рожденное для того, чтобы ползать, не является существом, которое может летать». Частноутвердительным оказывается суждение «некоторые грибы съедобны», а частноотрицательным — также

справедливое суждение «некоторые грибы не съедобны».

Пожалуй, наиболее просто, наглядно и вместе с тем строго силлогистику Аристотеля можно изложить с помощью кругов Эйлера, которые рассматривались выше.

В любом из четырех категорических суждений как левый термин *S*, так и правый термин *P* определяют некоторые классы предметов, и отношение между этими классами в общем случае можно показать пересекающимися кругами Эйлера (рис. 4.1). Напомним, что любые конкретные предметы изображаются на рисунке некоторыми точками, и если точка находится внутри круга, то это значит, что предмет принадлежит соответствующему классу. Рассматриваемые классы будем обозначать теми же буквами *S* и *P*.

На рис. 4.1 видно, как при пересечении круги Эйлера образуют четыре области, разбивая тем самым множество всех предметов на четыре части.

I. Предметы принадлежат классу *S*, но не принадлежат классу *P*.

II. Предметы принадлежат одновременно обоим классам.

III. Предметы принадлежат классу *P*, но не принадлежат классу *S*.

IV. Предметы не принадлежат ни классу *S*, ни классу *P*.

Заметим, что в силлогистике Аристотеля последний случай не исследуется. Каждое из конкретных категорических суждений сообщает некоторую информацию либо об области I, либо об области II (область III, как нетрудно видеть, симметрична области I).

Так, общеутвердительное суждение *A* (все *S* суть *P*)

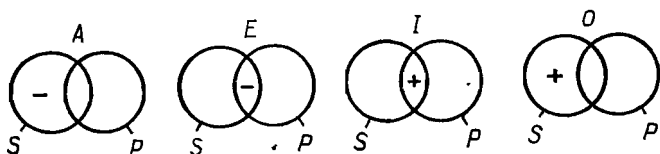


Рис. 4.2

утверждает, что область I является пустой. В теоретико-множественной символике это можно выразить формулой $S \cap P = \emptyset$. Общеотрицательное суждение *E* представляет собой утверждение о пустоте области II: $S \cap P = \emptyset$. Частноутвердительное суждение *I* (некоторые *S* суть *P*) говорит, что область II содержит по крайней мере один предмет, т. е. не является пустой ($S \cap P \neq \emptyset$), а частноотрицательное суждение *O* (некоторые *S* не суть *P*) говорит то же об области I ($S \cap P \neq \emptyset$).

Отмечая пустоту области знаком минус (—), а непустоту — знаком плюс (+), представим суждения в более наглядной форме (рис. 4.2).

Важной особенностью силлогистики Аристотеля является то, что она разработана для непустых классов (задаваемых терминами *S* и *P*), в то время как современные логические исчисления допускают, что некоторые из рассматриваемых классов могут оказаться пустыми. Именно поэтому частноутвердительное суждение *I* может рассматриваться как простое следствие общеутвердительного суждения *A*. Например, из того, что «все птицы» имеют перья, автоматически следует, что существует по крайней мере одна птица, которая имеет перья. Действительно, непустота области II вытекает из пустоты области I, поскольку известно, что класс *S* в целом не пуст: в нем содержится по крайней мере один объект и если он не может находиться в области I, то, значит, он принадлежит области II, так как класс *S* состоит только из этих двух областей. Аналогичным образом частноотрицательное суждение *O* вытекает из общеотрицательного суждения *E*.

Основное содержание силлогистики связано с пониманием силлогизма как умозаключения, в котором из двух категорических суждений, называемых посылками, получается третье, называемое заключением. Общее число терминов в суждениях равно при этом трем: один из них входит в обе посылки, называется средним и обозначается буквой *M* (*medius*), а два других, называемых крайними, переходят из посылок в заключение и обозначаются через *S* (меньший термин) и *P* (больший термин) в соответствии с тем местом, которое они займут в заключении: *S* будет играть роль субъекта, а *P* — роль предиката.

Что же касается посылок, то термины *M*, *S* и *P* могут занимать в них любое место, т. е. каждый из них может служить в посылке или субъектом, или предикатом. Посылка с большим термином *P* называется соответственно большей и ставится на первое место, посылка с меньшим термином *S* называется меньшей. Четыре возможных при этом варианта размещения терминов в посылках принято называть фигурами силлогизма:

Первая фигура	Вторая фигура	Третья фигура	Четвертая фигура
(<i>M</i> , <i>P</i>)	(<i>P</i> , <i>M</i>)	(<i>M</i> , <i>P</i>)	(<i>P</i> , <i>M</i>)
(<i>S</i> , <i>M</i>)	(<i>S</i> , <i>M</i>)	(<i>M</i> , <i>S</i>)	(<i>M</i> , <i>S</i>)

Выбирая в качестве посылок те или иные категорические суждения, можно заполнить каждую из этих фигур 16 способами, образуемыми комбинациями четырех вариантов выбора первого суждения с четырьмя вариантами выбора второго. Умножив это число на число фигур, получим 64 — число различных пар категорических суждений, связывающих крайние термины со средним. Однако не любая из этих пар порождает заключение, отличное от посылок, а лишь некоторые. Их оказывается 19, и именно они, вместе с вытекающими из них заключениями, называются силлогизмами. Найдём их.

Анализ отношений между тремя терминами *M*, *S* и *P* удобно проводить, рассматривая изображённую на рис. 4.3 фигуру, образованную тремя пересекающимися кругами Эйлера, подобно тому, как это делал Э. Беркли* в своем построении «силлогистической машины»,

* Беркли Э. Символическая логика и разумные машины. М., 1961.

демонстрируя возможность полной автоматизации силлогистических выводов.

Из рисунка видно, что заключение $A(S, P)$ — все S суть P — следует в тех случаях, когда оказываются пустыми области a и d . Но это возможно лишь в одном случае: когда справедливы суждения $A(M, P)$ и $A(S, M)$, как это показано на рис. 4.4, где заштрихованы области

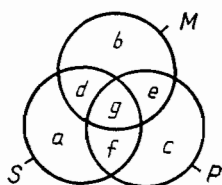


Рис. 4.3

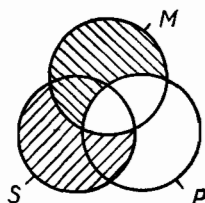


Рис. 4.4

запрета, соответствующие таким посылкам. Так находится силлогизм

$$\frac{A(M, P) \quad A(S, M)}{A(S, P)} .$$

Над чертой перечисляются посылки, а под чертой показано заключение. Этот силлогизм называется традиционно *Barbara* — последовательность гласных в данном слове указывает на типы входящих в силлогизм суждений, и какого-либо другого смысла это слово здесь не имеет. Размещение терминов в силлогизме *Barbara* соответствует первой фигуре, и его можно представить строчкой символов $(1, A, A) \rightarrow A$.

Посмотрим, когда возможно заключение $E(S, P)$ — ни одно S не есть P . Очевидно (см. рис. 4.3), что оно справедливо в тех случаях, когда пустыми будут области f и g , а это возможно лишь в двух ситуациях, показанных на рис. 4.5:

1) когда $A(P, M)$ и когда либо $E(S, M)$, либо $E(M, S)$;

2) когда $A(S, M)$ и когда либо $E(P, M)$, либо $E(M, P)$. Так находим еще четыре силлогизма:

$$(2, A, E) \rightarrow E(Camestres),$$

$$(4, A, E) \rightarrow E \text{ (Camenes)},$$

$$(1, E, A) \rightarrow E \text{ (Celarent)},$$

$$(2, E, A) \rightarrow E \text{ (Cesare)}.$$

Обратим внимание на симметричность отношения E , помогающую отыскивать силлогизмы. Выражения $E(S, M)$ и $E(M, S)$ равносильны, что легко усмотреть из пересекающихся кругов Эйлера, и поэтому если, к

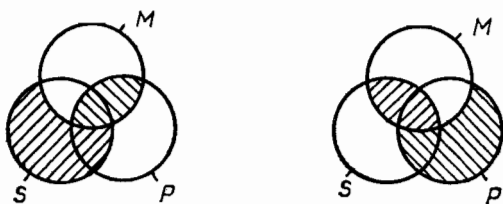


Рис. 4.5

примеру, некоторый вывод следует при $E(S, M)$, то он будет справедлив и при $E(M, S)$ при сохранении других условий. Аналогичной симметричностью обладает и отношение I .

Далее рассмотрим условия вывода заключения $I(S, P)$. Теперь надо доказать непустоту области, образуемой компонентами f и g . Очевидно, что она будет не пуста, если не пуста некоторая более широкая область, включающая f и g , и если будет показано, что остальные ее компоненты пусты.

Отмечая непустоту некоторой составной области, окаймляющей ее пунктирной линией (напомним, что каждая из областей M , S и P считается непустой), покажем возможные ситуации, в которых справедливо заключение $I(S, P)$ — некоторые S суть P .

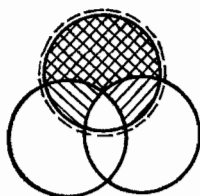


Рис. 4.6

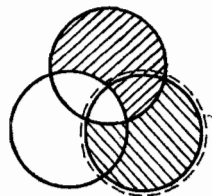


Рис. 4.7

Из рис. 4.6 следует силлогизм

$$(3, A, A) \rightarrow I(Darapti),$$

а из рис. 4.7 — силлогизм

$$(4, A, A) \rightarrow I(Bramalip).$$

Заметим, что рассматриваемое частноутвердительное заключение $I(S, P)$ следует и в симметричной ситуации $(1, A, A)$, но оно поглощается там общеутвердительным

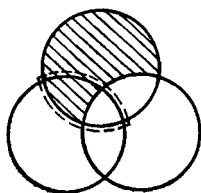


Рис. 4.8

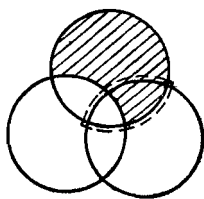


Рис. 4.9

заключением $A(S, P)$. Далее с помощью аналогичных рассуждений из рис. 4.8 находим силлогизмы

$$(1, A, I) \rightarrow I(Darii),$$

$$(3, A, I) \rightarrow I(Datisi),$$

а из рис. 4.9 — силлогизмы

$$(3, I, A) \rightarrow I(Disamin),$$

$$(4, I, A) \rightarrow I(Dimaris).$$

Наконец, найдем силлогизмы с заключением $O(S, P)$ — некоторые S не суть P . Для этого надо выявить ситуации, в которых оказывается непустой область, образованная компонентами a и d . Здесь возможны следующие

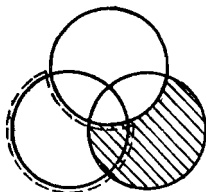


Рис. 4.10

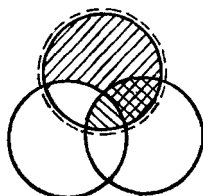


Рис. 4.11

варианты, получаемые с учетом симметричности отношений E и I .

Из рис. 4.10 следует силлогизм

$$(2, A, O) \rightarrow O \text{ (Baroco),}$$

из рис. 4.11 — силлогизмы

$$(3, E, A) \rightarrow O \text{ (Felapton),}$$

$$(4, E, A) \rightarrow O \text{ (Fesapo),}$$

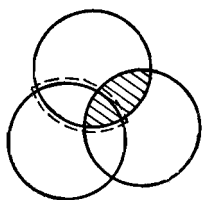


Рис. 4.12

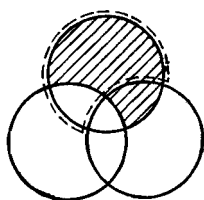


Рис. 4.13

из рис. 4.12 — четыре силлогизма

$$(1, E, I) \rightarrow O \text{ (Ferio),}$$

$$(2, E, I) \rightarrow O \text{ (Festino),}$$

$$(3, E, I) \rightarrow O \text{ (Ferison),}$$

$$(4, E, I) \rightarrow O \text{ (Fresison),}$$

из рис. 4.13 — последний силлогизм

$$(3, O, A) \rightarrow O \text{ (Bocardo).}$$

Итак, мы нашли 19 силлогизмов. Как это явствует из рисунков, в четырех из них учитывается предположение о непустоте некоторых из рассматриваемых классов: силлогизм *Bramalip* справедлив, если не пуст класс P , а силлогизмы *Darapti*, *Felapton* и *Fesapo* — если не пуст класс M . Поскольку современная логика не связана с таким предположением, в ней обычно принимаются лишь остальные 15 силлогизмов. Что же касается четырех упомянутых силлогизмов, то их также можно сформулировать в рамках современной логики, если добавить к двум посылкам третью, показывающую в явном виде

используемое предположение о непустоте класса P или M .

Для удобства представим полученные силлогизмы в виде следующей таблицы, столбцы которой соответствуют различным фигурам, а строки — различным комбинациям посылок, входящих в эти силлогизмы. Число таких комбинаций оказалось равным восьми.

Посылки	Фигуры				
		1	2	3	4
	1 2	$M-P$ $S-M$	$P-M$ $S-M$	$M-P$ $M-S$	$P-M$ $M-S$
1, 2					
(A, A)	A	—	I	I	
(A, E)	E	E	—	E	
(A, I)	I	—	I	—	
(A, O)	—	O	—	—	
(E, A)	—	E	O	O	
(E, I)	O	O	O	O	
(I, A)	—	—	I	I	
(O, A)	—	—	O	—	

Если учесть предположение о непустоте класса S , то к найденным 19 правильным силлогизмам можно добавить еще пять, производных от тех силлогизмов, в которых заключением служит суждение общего типа A или E — а таких силлогизмов ровно пять. Действительно, если класс S не пуст и если становится известно, что «все S суть P », то очевидна справедливость и такого утверждения: некоторые S суть P . Аналогично из суждения «ни одно S не есть P » логически следует «некоторые S не суть P », если, конечно, существуют объекты, принадлежащие классу S .

Выпишем эти пять дополнительных силлогизмов:

- (1, A, A) $\rightarrow$ I,
- (1, E, A) $\rightarrow$ O,
- (2, A, E) $\rightarrow$ O,
- (2, E, A) $\rightarrow$ O,
- (4, A, E) $\rightarrow$ O.

Силлогистику Аристотеля, выраженную в современ-

ной символике, можно интерпретировать как простейшую систему распознавания. При этом термины S , M и P трактуются как признаки, которыми могут обладать или не обладать рассматриваемые предметы, образующие некоторый универсальный класс. Суждения типа A и E представляют собой запреты на некоторые комбинации признаков: утверждается, что не существует объектов с такими комбинациями. Суждения типа O и I , напротив, говорят о существовании таких предметов.

Силлогизмы с заключениями типа A и E представляют собой умозаклучения, в которых из заданных двух запретов делается вывод об отличном от них третьем запрете. А все остальные силлогизмы можно свести к схеме, в которой в исходных посылках содержится утверждение о существовании некоторого предмета с частично заданными свойствами — для некоторых признаков сообщается, обладает ли ими этот предмет. На основании информации, содержащейся в других посылках, свойства предмета уточняются — определяется, обладает или нет рассматриваемый предмет некоторыми из остальных признаков. А это и есть задача распознавания! После такого уточнения делается вывод о том, что рассматриваемый предмет принадлежит или области $S \cap P$ (в этом случае выводится заключение типа I), или области $S \cap \bar{P}$ (здесь принимается заключение типа O).

Рассмотрим, к примеру, силлогизм *Darii*, часто иллюстрируемый следующим знаменитым рассуждением:

Все люди смертны.
Сократ—человек.
—————
Сократ—смертен.

Силлогизм содержит две посылки:

$A(M, P),$

$I(S, M).$

Вторая из них означает, что существует некоторый предмет, обладающий и признаком S , и признаком M . Но поскольку не существует предметов, обладающих признаком M и не обладающих при этом признаком P (как утверждает первая посылка), остается признать, что рассматриваемый предмет по необходимости обладает также признаком P . Поэтому он принадлежит области

$S \cap P$, которая, следовательно, не пуста. Так получается вывод:

$$I(S, P).$$

Аналогичные рассуждения могут проводиться и при рассмотрении так называемых полисиллогизмов, исследование которых было начато Аристотелем и продолжено Теофрастом и другими его учениками. Полисиллогизмом называется умозаключение с более чем двумя посылками, например:

$$\begin{array}{l} \text{Все } P \text{ суть } Q, \\ \text{Все } Q \text{ суть } R, \\ \text{Все } R \text{ суть } S, \\ \hline \text{Все } P \text{ суть } S. \end{array}$$

В общем случае число посылок в полисиллогизме может быть произвольным, но лишь одна из них может принадлежать типу I или O , т. е. оказаться суждением о существовании некоторого объекта. Поскольку число терминов, связываемых полисиллогизмом, может быть довольно большим, использование кругов Эйлера при поиске следующего из полисиллогизма заключения становится нецелесообразным — уж больно сложную картину образует пересечение этих кругов в общем случае. Да и сами «круги» пришлось бы так деформировать, что они уже не походили бы на круги в геометрическом смысле слова — ведь надо ради общности рассмотрения, чтобы n кругов делили плоскость на 2^n области, что возможно лишь при $n \leq 3$.

Воспользуемся, однако, тем обстоятельством, что каждая из посылок связывает лишь два термина, или признака. Это позволяет изобразить произвольную совокупность суждений в виде своеобразной сети, состоящей из узлов и звеньев, соединяющих некоторые узлы попарно. Ограничимся при этом пока рассмотрением суждений общего типа — суждений-запретов типа A или E , или суждений о несуществовании объектов с некоторыми свойствами. Узлы сети пусть представляют термины-признаки $a, b, \dots$, а звенья задают суждения, т. е. импликативные связи между признаками, именно те, которые мы раньше договорились представлять элементарными конъюнкциями второго ранга: $ab, \bar{a}c, \bar{c}d$ и т. д.

Связи уточняются с помощью белых и черных кружочков, помещаемых на концах звена и показывающих,

каким образом символ признака входит в элементарную конъюнкцию: если кружок белый, то символ входит под знаком отрицания, если черный — без знака отрицания (рис. 4.14).

Символическое представление связи	Графическое представление связи
ab	$a \bullet \text{-----} \bullet b$
$\bar{a}b$	$a \circ \text{-----} \bullet b$
$a\bar{b}$	$a \bullet \text{-----} \circ b$
$\bar{a}\bar{b}$	$a \circ \text{-----} \circ b$

Рис. 4.14

Пусть, например, задана совокупность посылок $A(a, b)$, $A(b, c)$, $E(b, d)$, $E(d, e)$, $A(c, e)$, $A(f, e)$, определяющая систему запретов на комбинации признаков. Эта система образует в пространстве признаков

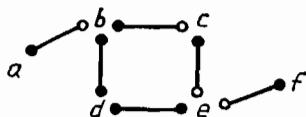


Рис. 4.15

область запрета, которую можно представить посредством ДНФ $a\bar{b} \vee b\bar{c} \vee b\bar{d} \vee d\bar{e} \vee c\bar{e} \vee f\bar{e}$. Изобразим ее графически, построив по описанным правилам сеть посылок (рис. 4.15). Спрашивается, какие логические выводы могут следовать из этой системы и как их получить?

Оказывается, что все выводы легко найти, глядя на сеть посылок. Справедливо следующее правило: если некоторые два узла сети соединены цепочкой, в которой любые соседние звенья стыкуются разноцветными концами (один — черный, другой — белый), то их можно соединить напрямую звеном с такими же концами, как и у цепочки. Введенное таким образом в сеть новое звено будет представлять суждение-следствие. Подчеркнем, что это правило является исчерпывающим, обеспечивая нахождение всех суждений-следствий.

Цепочки с отмеченным свойством будем называть проводящими. В рассматриваемом примере проводящие

цепочки соединяют узлы, образующие следующие пары (кроме тех, которые связаны непосредственно одним звеном): (a, c) , (a, e) , (a, d) , (b, e) и (d, f) . Введя в сеть соответствующие новые звенья, получаем сеть, изображенную на рис. 4.16. Добавленные при этом в сеть звенья и представляют полную систему суждений-следствий, а именно совокупность суждений $A(a, c)$, $A(a, e)$, $E(a, d)$, $A(b, e)$ и $E(d, f)$.

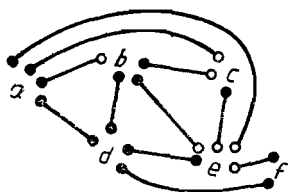


Рис. 4.16

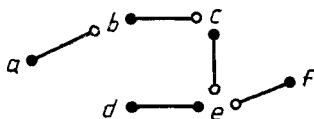


Рис. 4.17

Конечно, не обязательно загромождать рисунок сети, добавляя в нее новые звенья. Суждения-следствия можно выписывать и без этого, достаточно лишь обнаружить, что соответствующие узлы связаны некоторой проводящей сетью, и учесть окраску подходящих к узлам концов цепи.

Может оказаться, что при анализе некоторой сети посылок будут получены запреты-следствия вида $\bar{a}\bar{b}$ (не существует объектов, не обладающих ни признаком a , ни признаком b), не относящиеся к типу A или E , которые изучались Аристотелем. Однако в современной логике такие суждения занимают законное место наряду с другими.

Можно подойти к полисиллогизмам и с другой стороны, задавшись вопросом: а нельзя ли сократить заданную систему суждений, не нарушив содержащуюся в ней информацию? Действительно, некоторые из содержащихся в системе суждений могут оказаться логическими следствиями других, и в таком случае их удаление не повредит системе. Можно ли минимизировать в этом смысле систему посылок?

Очевидно, что эта задача эквивалентна рассмотренной ранее задаче минимизации ДНФ, задающей область запрета. Как уже отмечалось, задача минимизации ДНФ в общем случае довольно трудна, однако в данном случае нас спасает то обстоятельство, что ДНФ содержит конъюнкции

юнкции, ранг которых ограничен двумя. В связи с этим решение задачи существенно упрощается. Оказывается, что минимальная ДНФ единственна и находится совсем просто: любая последовательность допустимых упрощений — удалений некоторых элементарных конъюнкций — приводит к минимальной ДНФ.

В графическом варианте нахождение минимальной ДНФ, описывающей область запрета, выглядит следующим образом: из сети выбрасывается любое звено, если только оно порождается некоторой отличной от этого звена проводящей цепочкой. Действительно, его всегда можно будет получить в качестве логического вывода.

Так, в приведенном примере оказывается лишним звено, связывающее узлы b и d . Удалив его, получим сеть, из которой уже ничего нельзя выбросить: сеть оказывается тупиковой, или безызбыточной, причем единственной такого рода. Стало быть, она и минимальна. Результат показан на рис. 4.17.

Переходя от звеньев к соответствующим элементарным конъюнкциям, получаем минимальную ДНФ:

$$ab \vee bc \vee c\bar{e} \vee de \vee \bar{e}f.$$

Анализ системы суждений-посылок можно упростить еще более, если представить ее в виде перечня максимальных проводящих цепочек, т. е. таких, которые не содержатся целиком в каких-либо других проводящих цепочках. В данном случае таких цепочек две: одна из них соединяет узлы a и d , другая — узлы d и f . Эти цепочки состоят из следующих звеньев:

$$1) a \rightarrow b, b \rightarrow c, c \rightarrow e, e \rightarrow d,$$

$$2) d \rightarrow \bar{e}, \bar{e} \rightarrow \bar{f}$$

и их можно представить чисто условно выражениями

$$1) a \rightarrow b \rightarrow c \rightarrow e \rightarrow d,$$

$$2) d \rightarrow \bar{e}, \bar{e} \rightarrow \bar{f}$$

(не путая последние выражения с формулами композиций операции импликации).

Например, первая из цепочек означает следующее: если объект обладает признаком a , то он обладает признаком b ;

если объект обладает признаком b , то он обладает признаком c ;

если объект обладает признаком c , то он обладает признаком e ;

если объект обладает признаком e , то он не обладает признаком d .

Заметим, что каждая проводящая цепочка порождает две эквивалентные цепочки импликаций, получаемые в зависимости от того, с какого конца анализируемой цепочки в сети начинаются рассуждения. Так, цепочка импликаций $a \rightarrow b \rightarrow c \rightarrow e \rightarrow \bar{d}$ была получена при рассмотрении соответствующей проводящей цепочки в сети, начиная с узла a . Если же проанализировать проводящую цепочку с другого конца (d), то найдем цепочку импликаций $d \rightarrow \bar{e} \rightarrow \bar{c} \rightarrow \bar{b} \rightarrow \bar{a}$, представляющую обратный ход рассуждений:

если объект обладает признаком d , то он не обладает признаком e ;

если объект не обладает признаком e , то он не обладает признаком c ;

если объект не обладает признаком c , то он не обладает признаком b ;

если объект не обладает признаком b , то он не обладает признаком a .

Логический вывод становится при этом совершенно прозрачным: бесспорно, что из цепочки можно выбрать любую ее часть, содержащую не менее двух звеньев (иначе вывод будет тривиален, совпадая с уже имеющимся в системе суждением), и построить суждение-заклучение, соответствующее местоположению этой частичной цепочки и окраске ее концов. Очевидно, что есть смысл рассматривать цепочки, содержащие не менее двух звеньев, т. е. не менее трех узлов. При выполнении этого условия число различных нетривиальных заключений, порождаемых цепочкой с n узлами, равно, как нетрудно подсчитать, $n(n-3)/2+1$.

В приведенном примере первая цепочка порождает шесть заключений, вторая — только одно.

При рассмотрении произвольной системы суждений типа A и E , когда на нее не накладывается никаких ограничений, может оказаться, что она будет вырожденной. Это значит, что какой-то признак (а может быть, и несколько) будет несущественным, т. е. будет или присущ всем объектам, или никакому из них, или что в рассматриваемом классе вообще не существует никаких объектов — все комбинации признаков окажутся в этом случае под запретом.

Вырожденность сети легко выявляется при ее визу-

альном анализе. Если некоторая проводящая цепь обладает одноцветными концами и опирается на один и тот же узел, то это свидетельствует о несущественности соответствующего узлу признака — под запрет попадает его значение, представленное окраской концов цепи. В таком случае сеть можно упростить, удалив из нее все звенья, подходящие к данному узлу, и показав запрет одиночным кружком соответствующего цвета.

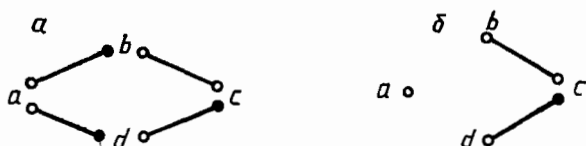


Рис. 4.18

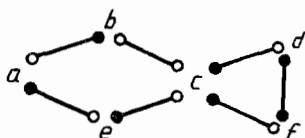


Рис. 4.19

Например, в сети, представленной на рис. 4.18,а, под запретом оказывается значение 0 признака *a*, другими словами, любой объект рассматриваемого класса должен обладать этим признаком. Упрощенная сеть показана на рис. 4.18,б.

Если под запретом окажутся оба значения некоторого признака, то это будет означать, что рассматриваемый класс пуст. Иллюстрацией к этому может служить сеть, изображенная на рис. 4.19. Критическим оказывается здесь признак *c* — на соответствующий узел опираются две проводящие цепи, запрещающие сообщать оба значения признака *c*.

А теперь перейдем к собственно задаче распознавания. Допустим, нам известно, что некоторый объект обладает определенным признаком (или известно, что он им не обладает), и требуется экстраполировать его свойства, вычислив, если возможно, значения остальных признаков. Вот для этого и удобно использовать сеть посылок, выполнив следующую операцию.

Найдем в сети узел, соответствующий заданному при-

знаку, и как бы потянем за кончики подходящих к нему звеньев, если окраска этих кончиков соответствует значению признака в рассматриваемом объекте. Предположим, что если к противоположным концам этих звеньев подходят иначе окрашенные концы некоторых других звеньев, то они также потянутся — соседние звенья слипаются разноцветными концами. В свою очередь новые звенья могут потянуть за собой еще какие-то из осталь-

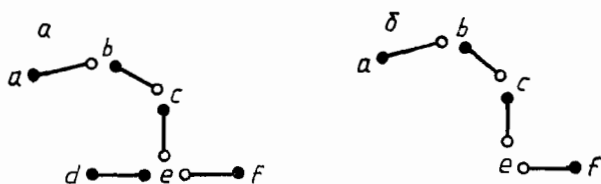


Рис. 4.20

ных и т. д. В результате из сети будет извлечена некоторая ее часть, состоящая из совокупности звеньев, определенным образом ориентированных относительно выбранного узла: один конец каждого звена окажется ближним, другой — дальним. Расцветка дальних концов покажет на запретные значения некоторых признаков для рассматриваемого объекта. Следовательно, объект должен обладать противоположными значениями этих признаков.

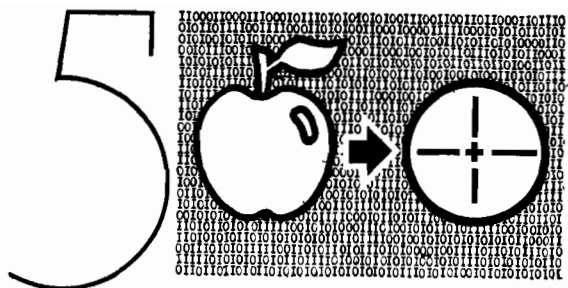
Так, если система посылок задана сетью, показанной на рис. 4.20, а, и если известно, что некоторый объект из рассматриваемого класса не обладает признаком e , то для уточнения свойств этого объекта надо «потянуть» за звенья, подходящие к узлу e светлыми концами. Тогда из сети будет извлечена часть (рис. 4.20, б), однозначно определяющая значения всех фигурирующих в ней признаков. Оказывается, что $f=0$, $c=0$, $b=0$ и $a=0$, т. е. данный объект не обладает ни одним из признаков f , c , b и a .

Так решается задача распознавания в том случае, когда известно значение лишь одного из признаков у рассматриваемого объекта. В случае, когда известны значения нескольких признаков, задача решается аналогично, и при этом каждый признак «тянет» за собой несколько новых, так что остается лишь объединить получаемые таким образом результаты.

Следует, однако, учитывать, что возможно возникно-

вление противоречия — когда оба значения некоторого признака окажутся под запретом. Такую ситуацию можно выявить следующим формальным приемом: если потянуть сеть сразу из всех узлов, соответствующих признакам с известными значениями, то обнаружится, что некоторые из них окажутся связанными цепочками со слипшимися звеньями. Другими словами, будет обнаружено наличие в сети такой проводящей цепи, которая соединяет некоторые два узла из числа указанных выше, причем окраска концов этой цепи будет указывать на то, что заданная комбинация значений признаков находится под запретом.

В таком случае придется, очевидно, допустить, что или выполненное ранее определение значений признаков у рассматриваемого объекта было неправильным, или ошибочной оказалась силлогистическая модель, или, наконец, рассматриваемый объект вообще не относится к тому классу, для которого эта модель адекватна.



РАСПОЗНАВАНИЕ КАК ЛОГИЧЕСКИЙ ВЫВОД

В традиционной постановке задачи распознавания один из признаков объявляется целевым, или признаком класса, принадлежность к которому требуется выяснять. Считается, что на этапе обучения этот признак измеряется наравне с остальными, но на этапе собственно распознавания он недоступен для непосредственного измерения или во всяком случае такое измерение крайне нежелательно.

Например, известный в медицинской практике признак «воспаление аппендикса» может быть установлен непосредственно лишь при операции — вскрытии брюшной полости. Очевидно, что такое непосредственное измерение слишком дорого и с точки зрения врача, и тем более пациента. Поэтому целесообразно заменить его распознаванием на основе информации, содержащейся в других признаках: болевых ощущениях в определенной области живота, данных некоторых анализов и т. п.

Весьма заманчиво было бы решить задачу распознавания «полностью», выразив целевой признак как некоторую функцию остальных, т. е. измеряемых, признаков (измерение понимается здесь в широком смысле — его результат не обязательно выражается числом). Тогда в любом случае

значение целевого признака определялось бы однозначно. Неудивительно поэтому, что большинство работ в области распознавания посвящено поискам именно таких функций, называемых решающими. Как же они находятся?

Получаемые на этапе обучения результаты представляют в многомерном пространстве измеряемых признаков (в число которых не входит целевой) в виде своеобразных «созвездий», или «облаков», образуемых объектами-точками с одинаковыми значениями целевого признака. Решающая функция имеет вид некоторого уравнения поверхности, разделяющей эти созвездия. Тогда распознавание сводится к определению координат наблюдаемого объекта и

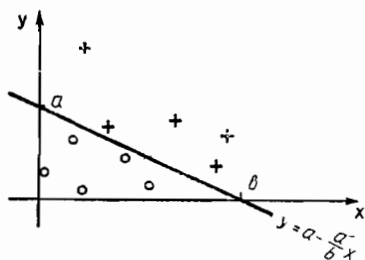


Рис. 5.1

выяснению того, по какую сторону разделяющей поверхности находится получаемая точка.

Такой подход особенно понятен, когда целевой признак двоичный, а измеряемые — многозначны. В простейшем случае, иллюстрируемом на рис. 5.1, роль разделителя может играть прямая, параметры которой a и b находятся на этапе обучения. Она разделяет «крестики» и «нолики» — входящих в обучающую выборку представителей двух классов. В данном примере решающую функцию удобно задать в виде неравенства $y > a - ax/b$. Если для некоторого нового объекта с координатами x и y это неравенство выполняется, то он относится к классу крестиков, если же не выполняется — к классу ноликов.

Если число измеряемых признаков равно не двум, а трем, прямая заменяется плоскостью, при большем числе признаков — гиперплоскостью; соответственно увеличивается и число параметров, определяемых при обучении. Такой метод распознавания называется линейным. Однако он не всегда применим. Может оказаться, что нельзя «остроить» плоскость, разделяющую представителей различных классов. Тогда прямые и плоскости заменяются ломаными (рис. 5.2) и поверхностями, составленными из кусочков плоскостей, или же более сложными кривы-

ми и поверхностями, описываемыми нелинейными уравнениями.

Однако и это не всегда возможно. Упомянутые созвездия могут переплетаться, настолько глубоко проникая друг в друга, что вообще не представляется возможным провести достаточно четкую границу, разделяющую их.

Да и в том, казалось бы, простейшем случае, когда

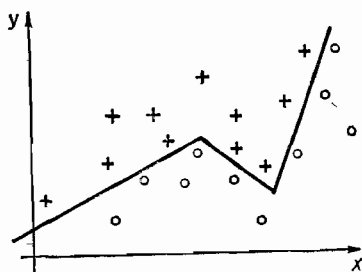


Рис. 5.2

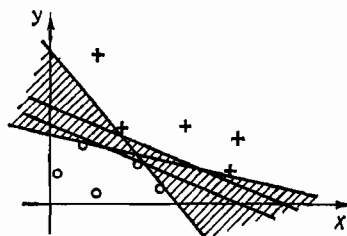


Рис. 5.3

удается разделить созвездия плоскостью, возникает ряд совсем не простых вопросов. Ведь, как правило, провести разделяющую плоскость можно по-разному, как это показано на рис. 5.3 (с теми же данными, что и на рис. 5.1). В результате образуется некоторая область неопределенности, элементы которой могут быть отнесены как к одному, так и к другому классу в зависимости от выбора плоскости. Похоже, что лучше отказаться от распознавания таких объектов, чем пытаться определять для них значение целевого признака.

Поэтому вместо решающей функции удобнее использовать решающее правило, предусматривающее возможность появления критических ситуаций. Если все признаки двоичны, то его можно представить как совокупность двух булевых функций φ_1 и φ_0 от измеряемых признаков: функция φ_1 определяет условия, при которых целевой признак может принять лишь значение 1, функция φ_0 задает аналогичные условия для значения 0. Другими словами, эти функции служат импликантами: если $\varphi_1=1$, то анализируемый объект обладает целевым признаком, если $\varphi_0=1$, то не обладает. В то же время из нулевых значений φ_1 и φ_0 ничего не следует.

Отличие решающего правила от решающей функции

заключается в том, что при его применении может оказаться, что обе функции φ_1 и φ_0 примут одновременно значение 0 или значение 1. В первом случае это будет говорить о том, что значение целевого признака остается неопределенным, т. е. может быть любым, а во втором — что наблюдаемый объект противоречит используемой при распознавании модели класса, заданной в виде совокупности импликативных закономерностей.

Для силлогистической модели мира функции φ_1 и φ_0 находятся довольно просто. Рассмотрим, к примеру, модель, представленную на рис. 4.20. Выбрав в качестве целевого признака c , легко найти, что этот признак принимает значение 1, когда $a=1$ или $b=1$, и значение 0, когда $d=1$ или $e=0$, или $f=1$. Отсюда следует, что $\varphi_1 = a \vee b$, а $\varphi_0 = d \vee \bar{e} \vee f$. Если обе функции примут значение 0 (а это возможно только при $(a, b, d, e, f) = (0, 0, 1, 0)$), то остается неясным, обладает наблюдаемый объект признаком c или нет. Если же обе функции обратятся в единицу (например, при $(a, b, d, e, f) = (1, 0, 0, 0, 0)$), наблюдаемый объект не укладывается в рамки заданной модели, в которой показанная комбинация значений признаков запрещена.

Итак, аппарат силлогистики позволяет решать задачи распознавания объектов с частично определенными свойствами логическим путем: на базе выявленных на этапе обучения импликативных закономерностей второго ранга. Это — весьма сильные закономерности, легко обнаруживаемые на обучающих выборках сравнительно небольшого размера. При более широком эксперименте можно выявлять и более тонкие закономерности, представляемые импликативными закономерностями более высоких рангов. Покажем, что и в данном случае задачу распознавания можно решать чисто логическим путем, представляя процесс ее решения в виде последовательности логических рассуждений, приводящих к тому или иному заключению. Решающий процесс можно рассматривать при этом как логический вывод, аналогичный тому, который проводится при автоматическом доказательстве теорем.

Посмотрим, как он строится в общем случае, когда ранги импликативных закономерностей не ограничены цифрой 2. Обратимся к примеру.

Допустим, что предметами распознавания служат представители некоторого класса объектов, моделиру-

емых точками в булевом пространстве признаков a, b, c, d, e, f . Предположим также, что на этапе обучения наблюдению подверглось 70 конкретных объектов, в результате чего составлена следующая таблица, в которой некоторые строки могут обладать одинаковыми значениями:

	a	b	c	d	e	f
Объект 1	1	1	0	1	0	0
Объект 2	0	1	1	1	1	1
...						
Объект 70	0	0	1	0	0	0

Таблица полностью не приводится, чтобы не загромождать текст.

В рассматриваемом примере при построении модели класса можно обратить внимание на оказавшиеся пустыми интервалы третьего ранга и выдвинуть гипотезы о соответствующих импликативных закономерностях. В самом деле, достоверность гипотез такого рода тем выше, чем меньше вероятность того, что интервал оказался пустым чисто случайно, а последняя оценивается сверху, как это было показано в главе 3, величиной $W(n, m, l)$, в нашем случае равной $W(6, 70, 3) = 0,014$. Это значение можно с некоторой натяжкой счесть достаточно малым, чего никак нельзя сказать, например, о значении $W(6, 70, 4) = 2,62$, оценивающем сверху вероятность аналогичного события для интервалов ранга 4.

Предположим, что среди интервалов, ранги которых не превышают трех, пустыми оказались лишь те, которые представлены строками следующей троичной матрицы:

$$T = \begin{bmatrix} a & b & c & d & e & f \\ 1 & - & 1 & - & - & 0 \\ - & 1 & - & - & 0 & 1 \\ 0 & - & - & 0 & 1 & - \\ - & 0 & - & 1 & - & 1 \\ - & 0 & 0 & 0 & - & - \end{bmatrix}$$

Именно эта матрица и рассматривается далее как модель исследуемого класса. Теперь ее строки интерпретируются как импликативные закономерности: например, первая строка утверждает, что в данном классе не существует объектов, обладающих признаками a и c , но не облада-

ющих в то же время признаком f . Данную модель можно задать и в алгебраической форме — посредством ДНФ запрета

$$\varphi = ac\bar{f} \vee be\bar{f} \vee \bar{a}\bar{d}e \vee \bar{b}df \vee \bar{b}\bar{c}\bar{d}.$$

Чтобы представить содержащуюся здесь информацию более наглядно, изобразим функцию запрета на геометрической модели булева пространства M , известной под

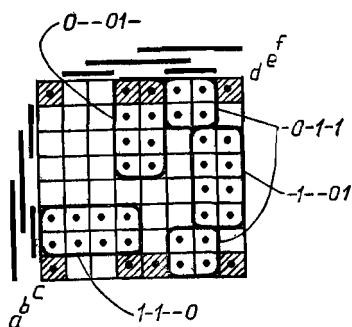


Рис. 5.4

названием карты Карно. Это — двумерная таблица, в которой каждый элемент пространства, т. е. каждая конкретная комбинация значений признаков, представляется отдельной клеткой. Соответствие между клетками и элементами задается обычно с помощью жирных линий, проводимых слева и сверху таблицы: линии помечаются символами булевых переменных и свидетельствуют о том, что располо-

женные против них (справа или внизу) клетки изображают элементы булева пространства с единичными значениями этих переменных. При этом элементы, на которых изображаемая функция принимает значение 1, как-то отмечаются, например точками. Очевидно, что при изображении булевой функции от n переменных потребуется таблица с 2^n клетками.

В таком виде рассматриваемая функция запрета φ показана на рис. 5.4. Например, левая верхняя клетка задает элемент 0 0 0 0 0 0 булева пространства переменных a, b, c, d, e, f , соседняя с ней справа — элемент 0 0 0 1 0 0, а соседняя снизу — элемент 0 0 1 0 0 0. Как легко увидеть, при движении по таблице сверху вниз изменяются значения переменных a, b, c , а при движении слева направо — значения переменных d, e, f . Заметим, что соседние элементы отличаются значением только одной переменной. Это объясняется тем, что при построении карты Карно используется симметричный код Грея, который каждому натуральному числу, начиная с нуля, ставит в соответствие свою кодовую комбинацию — булев вектор-константу, число компонент в кото-

ром выбирается в зависимости от количества кодируемых чисел:

0	0	0	0	0
1	0	0	0	1
2	0	0	1	1
3	0	0	1	0
4	0	1	1	0
5	0	1	1	1
6	0	1	0	1
7	0	1	0	0
8	1	1	0	0
...				

Код Грея удобен тем, что он приводит к легко обозримому представлению интервалов, в чем можно убедиться, посмотрев на рис. 5.4. В рассматриваемом примере все интервалы представлены довольно компактными областями таблицы, лишь элементы интервала — 0 0 0 — — оказались рассеянными, тем не менее они образуют легко воспринимаемую симметричную картину (эти элементы отмечены штриховкой).

Как же можно использовать представленную на карте Карно информацию при анализе нового объекта, у которого «измерены» лишь некоторые признаки? Очевидно, следует иметь в виду, что область допустимых (в предположении, что данная модель справедлива) объектов задается дополнением области запрета, а именно совокупностью неотмеченных элементов таблицы.

Допустим, удалось установить, что наблюдаемый объект обладает признаком a , но не обладает признаком f . Что можно сказать в этом случае о других признаках, не прибегая к их непосредственному измерению?

Обратимся к карте Карно. Очевидно, при $a=1$ и $f=0$ объект может попасть лишь в левый нижний квадрант карты, в связи с чем можно ограничиться рассмотрением именно этого квадранта, задающего интервал 1 — — — 0. Вырезав его, получим карту Карно уже не для шести переменных, а для четырех — удалены переменные с заданными значениями (рис. 5.5). Здесь, как и прежде, точками отмечены запретные клетки, и объект может попасть лишь в одну из оставшихся. Легко видеть, что при этом объект никак не может обладать признаком c и должен обязательно обладать по крайней мере одним

из признаков b или d . Что касается признака e , то он оказывается не связанным с результатами наблюдений, так что неизвестно, обладает им объект или нет — картина полностью симметрична относительно центральной вертикальной оси, на которой меняет значение признак e .

Тот же результат можно получить и более формальным путем, покрыв область запрета парой интервалов (как показано на рисунке) и выписав соответствующую

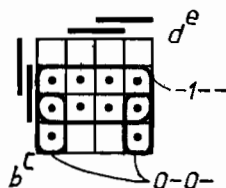


Рис. 5.5

ДНФ запрета $c \vee b\bar{d}$. Объект должен обладать такими значениями признаков, чтобы эта ДНФ обратилась в нуль, а это значит, что $c=0$ и $b=1$ или $d=1$ (допускается и одновременное принятие значения 1 признаками b и d).

Так делается полный прогноз по результатам наблюдений, и очевидно, что к нему ничего нельзя добавить.

В случае, когда число признаков велико, пользоваться картой Карно трудно или даже практически невозможно. Например, при 20 признаках потребовалась бы карта Карно, содержащая 2^{20} , т. е. свыше 1 000 000 клеток! В связи с этим полезно уметь решать рассматриваемую задачу чисто алгебраически.

Для этого достаточно подставить в исходную ДНФ запрета

$$\varphi = ac\bar{f} \vee be\bar{f} \vee \bar{a}\bar{d}e \vee \bar{b}d\bar{f} \vee \bar{b}\bar{c}\bar{d}$$

значения $a=1$ и $f=0$, преобразовав функцию φ к частному виду, соответствующему именно этим значениям:

$$\begin{aligned} \varphi(a=1, f=0) &= 1 \cdot c \cdot 1 \vee b \cdot \bar{e} \cdot 0 \vee 0 \cdot \bar{d} \cdot e \vee \bar{b} \cdot d \cdot 0 \vee \bar{b} \cdot \bar{c} \cdot \bar{d} = \\ &= c \vee \bar{b}\bar{c}\bar{d} = c \vee \bar{b}\bar{d}. \end{aligned}$$

В справедливости последнего равенства читатель может убедиться самостоятельно, проверив его на всех восьми

комбинациях значений фигурирующих в нем переменных.

Этот же результат можно получить и из матричного представления функции запрета — из троичной матрицы T , удалив в ней столбцы a и f и строки, в которых стоит 0 в столбце a или 1 в столбце f : очевидно, что представленные в таких строках запреты оказываются излишними в рассматриваемой ситуации. В результате такого упрощения матрицы T получается ее минор

$$\begin{array}{cccc} b & c & d & e \\ \left[\begin{array}{cccc} - & 1 & - & - \\ 0 & 0 & 0 & - \end{array} \right], \end{array}$$

которому соответствует полученная выше ДНФ запрета

$$c \vee \bar{b} \bar{c} \bar{d} = c \vee \bar{b} \bar{d}.$$

Область возможных значений признаков b , c , d и e для наблюдаемого объекта описывается функцией, инверсной данной:

$$\neg (c \vee \bar{b} \bar{d}) = \bar{c} (b \vee d).$$

Вернемся к классической постановке задачи распознавания, выделив некоторый целевой признак и считая, что при наблюдении новых объектов экспериментально определяются значения всех признаков, кроме целевого. Значение последнего требуется вычислить по решающему правилу, которое предварительно надо построить исходя из заданной системы закономерностей.

Пусть в данном примере роль целевого играет признак f . Обратимся к карте Карно, на которой представлена область запрета в булевом пространстве M — для удобства она заштрихована (рис. 5.6, a). Разобьем мысленно эту карту на две равные части, соответствующие различным значениям целевого признака f . Как видно, разрез пришелся на центральную вертикальную линию, на которой меняет свое значение переменная f . При этом на левой половине оказалась заданная область запрета нулевых значений признака f , а на правой — единичных. Перенесем эти области на общую карту Карно, представляющую булево пространство пяти признаков, остающихся после исключения f . Такой перенос можно осуществить, как бы сложив исходную карту по линии разреза, показанной на рисунке пунктиром, и удалив верхнюю

жирную линию, задающую значение признака f . Результат показан на рис. 5.6,б — после переворота правой половинки исходной карты направление штриховки соответственно изменилось, что позволяет отличать ее от штриховки левой части.

Таким образом, в булевом пространстве признаков a, b, c, d, e оказались выделенными две области: область M_1 , отмеченная прямой штриховкой, такой же, что исполь-

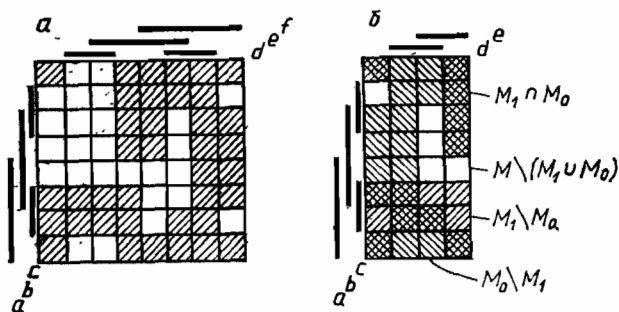


Рис. 5.6

зовалась на рис. 5.6,а, и область M_0 , заданная обратной штриховкой, полученной в результате переворота правой половинки исходной карты Карно. Область M_1 представляет собой область запрета значения 0 признака f , область M_0 — запрета значения 1 этого же признака.

Если результаты измерений признаков a, b, c, d, e таковы, что объект попадает в область M_1 , то это значит, что он может обладать только значением 1 целевого признака f , если же объект попадает в область M_0 , то в соответствии с используемыми закономерностями объект обладает значением 0 признака f . Другими словами, области M_1 и M_0 — не что иное, как характеристические множества булевых функций φ_1 и φ_0 , введенных в начале главы: множество M_1 образовано из элементов булева пространства, на которых принимает значение 1 функция φ_1 , а множество M_0 находится в таком же отношении с функцией φ_0 .

Однако, как видно из рис. 5.6,б, пространство переменных a, b, c, d, e оказалось разбитым не на две, а на четыре части:

- 1) только с прямой штриховкой ($M_1 \setminus M_0$),
- 2) только с обратной штриховкой ($M_0 \setminus M_1$),

3) одновременно с прямой и обратной штриховкой ($M_1 \cap M_0$),

4) без штриховки ($M \setminus (M_1 \cup M_0)$).

Чтобы провести полный анализ возможных при распознавании ситуаций, надо разобраться со смыслом этих областей.

Смысл первых двух областей по существу уже ясен: при попадании наблюдаемого объекта в область $M_1 \setminus M_0$ (когда выполняется условие $\varphi_1 \wedge \bar{\varphi}_0 = 1$) решающее правило должно приписывать признаку $\bar{f}$ значение 1, а при попадании в область $M_0 \setminus M_1$ (когда $\varphi_1 \wedge \varphi_0 = 1$) — значение 0.

Если же объект попадает в область $M_1 \cap M_0$ (когда $\varphi_1 \wedge \varphi_0 = 1$), придется признать, что оба значения 1 и 0 признака $\bar{f}$ для него оказываются запретными. Такой случай следует рассматривать как возникновение противоречия: результаты наблюдения за новым объектом явно противоречат построенной из импликативных закономерностей модели исследуемого класса.

Наконец, при попадании объекта в область $M \setminus (M_1 \cup M_0)$, когда $\varphi_1 \vee \varphi_0 = 0$, нельзя сделать никаких выводов, кроме одного, а именно признать, что информация, заключенная в модели и в наблюдениях, недостаточна для того, чтобы решить, обладает рассматриваемый объект признаком $\bar{f}$ или нет.

Анализ двух последних случаев выгодно отличает изложенный подход к построению решающего правила от известных методов «детерминированного» распознавания, при которых в любом случае наблюдаемый объект относится к одному из двух классов: считается, что объект или обладает признаком $\bar{f}$, или не обладает им. Очевидно, при такой обязательной однозначности приходится иногда делать определенные выводы без достаточных на то оснований. Можно вспомнить в связи с этим прогноз погоды, нередко подводящий легковых граждан, собирающихся на загородную прогулку в апреле, не захватив с собой плаща или зонтика...

До сих пор мы придерживались принципов так называемого «статического распознавания», при котором распознавание ведется на основе предварительного измерения всех признаков, кроме целевого. Однако эта информация может оказаться избыточной, в чем можно убедиться на следующем примере.

Пусть при наблюдениях за некоторым новым объек-

том того же класса, что и раньше, обнаружено, что объект обладает признаком d , но не обладает признаками a и e (т. е. $a=0, d=1, e=0$), значения же признаков b и c остаются неизвестными. Это означает, что объект принадлежит четырехэлементному интервалу, показанному на рис. 5.7,а. Сравнивая этот рисунок с предыдущим, легко увидеть, что данный интервал полностью содержится в области $M_0 \setminus M_1$. А отсюда следует, что

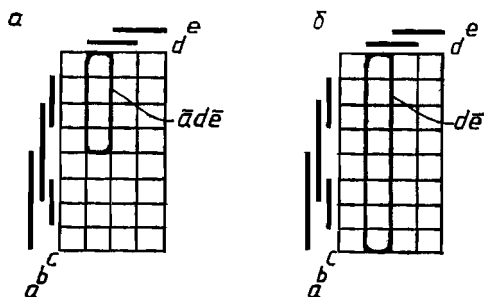


Рис. 5.7

каковы бы ни были значения признаков b и c , значения признака f в наблюдаемом объекте должно быть равно нулю.

Такой же вывод можно сделать и в том случае, когда будут известны значения лишь двух признаков d и e : $d=1, e=0$. Действительно, соответствующий интервал, показанный на рис. 5.7,б, пересекается только с двумя областями: $M_0 \setminus M_1$ и $M_1 \cap M_0$. Но область $M_1 \cap M_0$ интерпретируется как область невозможного, поэтому, пользуясь заданной системой закономерностей, мы полагаем, что рассматриваемый объект не может в нее попасть. Значит, он относится только к области $M_0 \setminus M_1$, откуда в свою очередь следует, что $f=0$.

Итак, распознавание ведется путем сравнения интервала возможного существования объекта, определяемого значениями измеренных признаков, с четырьмя областями, на которые разбито булево пространство признаков, доступных измерению. Однако, полагаясь на непогрешимость представленных областью запрета закономерностей, а также на безошибочность выполненных над объектом наблюдений, мы можем при распознавании

игнорировать возможность попадания объекта в область $M_1 \cap M_0$, ограничиваясь выявлением трех ситуаций:

- 1) интервал возможного существования объекта содержится целиком в области M_1 , в этом случае решающее правило предписывает целевому признаку значение 1;
- 2) интервал содержится в области M_0 , в этом случае целевому признаку предписывается значение 0;
- 3) в противном случае, когда интервал не содержится

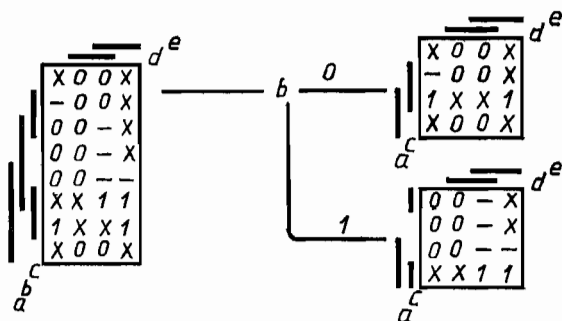


Рис. 5.8

целиком ни в M_1 , ни в M_0 , значение целевого признака остается неопределенным.

Такой подход особенно плодотворен при «динамическом распознавании», когда признаки наблюдаемого объекта измеряются последовательно. При этом преследуется цель поскорее перейти к одной из первых двух ситуаций, в которых значение целевого признака будет вычислено однозначно. При динамическом распознавании выбор очередного измеряемого признака ставится в зависимость от результатов измерения предыдущих.

Посмотрим, как это делается, все на том же примере, преобразовав заданную на рис. 5.6,б информацию к более удобному виду (рис. 5.8). Здесь на клетках карты Карно непосредственно показано, какие значения должны предписываться целевому признаку f на различных элементах булева пространства измеряемых признаков. Наряду с конкретными значениями 1 и 0 использованы символ неопределенности «—» и символ X, отмечающий запрещенные элементы, в которые, как мы считаем, наблюдаемый объект никогда не сможет попасть. Очевидно, что область M_1 состоит из элементов, отмеченных симво-

лами 1 или $\times$, а область M_0 — из элементов, отмеченных символами 0 или $\times$.

Пусть, например, первым измеряется признак b . При $b=0$ становится известно, что наблюдаемый объект относится к одной половине булева пространства, при $b=1$ — к другой. Очевидно, что соответствующую половину и стоит рассматривать в дальнейшем. Так производится расщепление карты Карно, показанное на

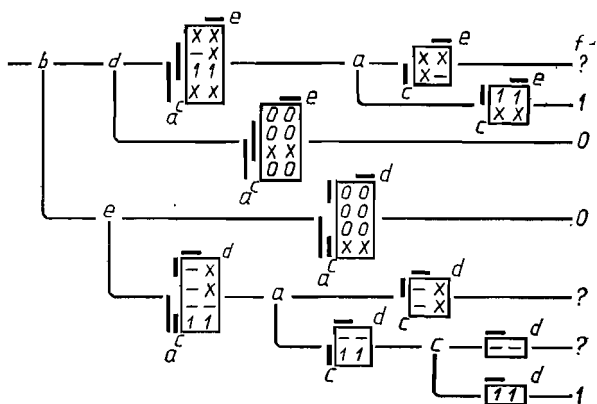


Рис. 5.9

рис. 5.8, где задана исходная карта, указаны оба значения признака b , по которому производится расщепление, и соответствующие им области последующего анализа.

Рассмотрим сначала вариант, в котором $b=0$. Какой следующий признак удобно выбрать для измерения? Пожалуй, признак d : если $d=1$, то ясно, что признак f должен иметь значение 0, если же $d=0$, то ситуация прояснится при дальнейшем измерении признака a — при $a=1$ признак f получит значение 1, при $a=0$ окажется, что значение признака f не может быть однозначно определено без непосредственного его измерения.

Эти и аналогичные рассуждения отображены на рис. 5.9 в форме своеобразного дерева, точки ветвления в котором отмечены символами подлежащих измерению признаков, а исходящие из этих точек ветви соответствуют значениям признаков: верхняя — значению 0, нижняя — 1. На ветвях изображены анализируемые участки карты Карно, а в крайней справа колонке представлены делаемые при прохождении соответствующих ветвей

выводы о значении целевого признака f : ему приписывается значение 0 или 1 или же оно остается неопределенным, что отмечается символом «—».

Показанные на дереве участки карты Карно играют здесь вспомогательную роль, поясняя выбор следующего признака для измерения. Если убрать эту информацию, то рисунок упростится — получим дерево распознавания признака f , изображенное на рис. 5.10. Оно представляет собой наглядную форму алгоритма распознавания, который можно выразить и словами, например следующим образом.

1. Измерить признак b . Если $b=1$, перейти к п. 2, если же $b=0$, измерить признак d . Если $d=1$, приписать признаку f значение 0 и выйти на п. 3. Если же $d=0$, то измерить признак

a . Если $a=1$, приписать признаку f значение 1 и выйти на п. 3. Если же $a=0$, отказаться от распознавания, выйдя на п. 3.

2. Измерить признак e . Если $e=0$, приписать признаку f значение 0 и выйти на п. 3. Если же $e=1$, измерить признак a . Если $a=0$, отказаться от распознавания, выйдя на п. 3. Если же $a=1$, измерить признак c . Если $c=0$, отказаться от распознавания, выйдя на п. 3. Если же $c=1$, приписать признаку f значение 1 и выйти на п. 3.

3. Конец алгоритма.

Конечно, описанный алгоритм динамического распознавания признака f не единственно возможный. Очевидно, что процесс распознавания можно начать с измерения какого-либо другого признака, отличного от b , тогда распознавание пойдет по другому пути.

При построении самого алгоритма распознавания не обязательно пользоваться картой Карно, тем более что при большом числе признаков это практически невозможно. Альтернативой является алгебраический подход, сводящий распознавание к решению логических уравнений.

При этом подходе прежде всего меняется представление — рассматриваемые теоретико-множественные кон-

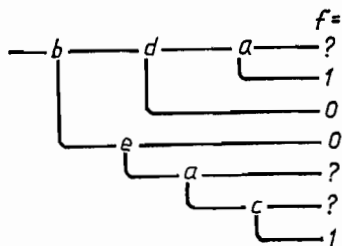


Рис. 5.10

фигурации, изображаемые на картах Карно, заменяются соответствующими алгебраическими выражениями.

Так, вместо области запрета, задаваемой на исходной карте Карно, рассматривается ее характеристическая функция (принимаяющая значение 1 на этой области) — булева функция запрета φ , а вместо областей M_1 и M_0 , где запрещены соответственно значения 0 и 1 целевого признака f — характеристические функции этих областей φ_1 и φ_0 , получаемые подстановкой в функцию φ указанных значений аргумента f . Для нашего примера

$$\varphi = ac\bar{f} \vee b\bar{e}f \vee \bar{a}\bar{d}e \vee bdf \vee \bar{b}\bar{c}\bar{d},$$

$$\varphi_1 = \varphi(f = 0) = ac\sqrt{a}\bar{d}e \vee \bar{b}\bar{c}\bar{d},$$

$$\varphi_0 = \varphi(f = 1) = b\bar{e} \vee \bar{a}\bar{d}e \vee b\bar{d} \vee \bar{b}\bar{c}\bar{d}.$$

Что касается интервала возможного существования объекта, то его алгебраическим представлением служит соответствующая элементарная конъюнкция k , образуемая из символов измеренных признаков. Например, если при наблюдении установлено, что объект обладает признаком a и не обладает признаком c , то $k = a\bar{c}$.

Утверждение, что наблюдаемый объект должен подчиняться закономерностям, т. е. не должен попасть в запретную для него область, принимает форму логического уравнения

$$\varphi = 0,$$

а результаты выполненных наблюдений оформляются в виде уравнения

$$k = 1,$$

говорящего о том, что объект находится в области, характеризующейся элементарной конъюнкцией k . Совокупность этих двух утверждений представляется уравнением

$$\bar{\varphi}k = 1,$$

которое и требуется решать в определенном смысле, а именно разрешать его относительно целевого признака f .

Как же оно решается?

Как и ранее, будем полагать, что наблюдаемый объект не противоречит закономерностям. Это значит, что

всегда должно существовать некоторое решение уравнения

$$\bar{\varphi}k = 1,$$

или, как говорят, формула $\bar{\varphi}k$ должна выполняться — должен существовать такой набор значений аргументов, который обращает ее в единицу.

Рассуждая аналогично тому, как это делалось при использовании карт Карно, мы вынуждены признать, что целевой признак f должен получить значение 1, если он не может получить значение 0, т. е. если будет тождественно уравнение

$$\bar{\varphi}k\bar{f} = 0,$$

и признак $\bar{f}$ должен получить значение 0, если он не может получить значение 1, т. е. если будет тождественно уравнение

$$\bar{\varphi}kf = 0.$$

Таким образом, дело сводится к проверке тождественности этих уравнений, которую можно выполнить следующим образом.

Ясно, что выражение $\bar{\varphi}k\bar{f}$ будет тождественно равно нулю, если функция φ будет принимать значение 1 всегда, когда принимает значение 1 конъюнкция $k\bar{f}$. А такую проверку легко осуществить непосредственно на матрице T , задающей функцию запрета φ . Для этого достаточно удалить из нее строки, ортогональные конъюнкции $k\bar{f}$, т. е. обращающие эту конъюнкцию в нуль (подстановкой указанных в строке значений аргументов), а также удалить столбцы, соответствующие аргументам k и $\bar{f}$. Остается посмотреть, будет ли функция, представленная полученным таким образом минором, тождественно равна единице. В случае положительного ответа на этот вопрос целевому признаку можно приписать значение 1.

Аналогично определяется возможность присвоения признаку f значения 0, только вместо конъюнкции $k\bar{f}$ рассматривается kf .

Процесс динамического распознавания, если он удачно спланирован, можно представить как такой порядок выбора измеряемых признаков, благодаря которому скорее обнаруживаются ситуации, где выполняется одно из указанных тождеств. Очевидно, что только при этом

может быть достигнута какая-то экономия в числе измеряемых признаков. Отсюда следуют естественные требования к способу построения алгоритма распознавания.

Однако вернемся к рассмотрению прежнего примера, построив для него алгоритм распознавания признака f уже не по карте Карно, а по матрице импликативных закономерностей T , пронумеровав для удобства ее строки:

$$T = \begin{array}{c} \begin{array}{cccccc} a & b & c & d & e & f \end{array} \\ \left[\begin{array}{cccccc} 1 & - & 1 & - & - & 0 \\ - & 1 & - & - & 0 & 1 \\ 0 & - & - & 0 & 1 & - \\ - & 0 & - & 1 & - & 1 \\ - & 0 & 0 & 0 & - & - \end{array} \right] \begin{array}{l} 1 \\ 2 \\ 3 \\ 4 \\ 5 \end{array} \end{array}$$

Есть смысл начинать распознавание с измерения признака, которому соответствует столбец, содержащий возможно большее число конкретных значений, поскольку в данном случае матрица будет сильнее упрощаться. Исходя из этого соображения, начнем с признака b .

При $b=0$ становится излишней строка 2, так как задаваемая ею область запрета не содержит элементов с таким значением признака b и, следовательно, не пересекается с интервалом возможного существования объекта, не обладающего признаком b . Удалив ее вместе со столбцом b , получим остаток

$$\begin{array}{c} \begin{array}{cccc} a & c & d & e & f \end{array} \\ \left[\begin{array}{cccc} 1 & 1 & - & - & 0 \\ 0 & - & 0 & 1 & - \\ - & - & 1 & - & 1 \\ - & 0 & 0 & - & - \end{array} \right] \begin{array}{l} 1 \\ 3 \\ 4 \\ 5 \end{array} \end{array}$$

Если же окажется, что $b=1$, далее следует анализировать такой остаток матрицы T :

$$\begin{array}{c} \begin{array}{cccc} a & c & d & e & f \end{array} \\ \left[\begin{array}{cccc} 1 & 1 & - & - & 0 \\ - & - & - & 0 & 1 \\ 0 & - & 0 & 1 & - \end{array} \right] \begin{array}{l} 1 \\ 2 \\ 3 \end{array} \end{array}$$

Рассмотрим первый вариант. В нем вторым измеря-

ется признак d . При $d=0$ матрица сокращается до

$$\begin{array}{cccc|c} a & c & e & f & \\ \hline [1 & 1 & - & 0] & 1 \\ 0 & - & 1 & - & 3 \\ - & 0 & - & - & 5 \end{array}$$

и далее измеряется признак a . При $a=0$ из матрицы удаляется единственная оставшаяся строка, в которой признак f имеет определенное значение. Очевидно, после этого никакие измерения оставшихся признаков, кроме непосредственно целевого, не могут привести к однозначному определению значения признака f . Если же $a=1$, далее будет анализироваться следующий остаток матрицы:

$$\begin{array}{ccc|c} c & e & f & \\ \hline [1 & - & 0] & 1 \\ 0 & - & - & 5 \end{array}$$

По нему видно, что значение 0 признака c запрещено. А при $c=1$ под запретом оказывается значение 0 целевого признака f . Стало быть, $f=1$.

Аналогично рассматривается второй вариант, соответствующий значению 1 признака b . В данном случае вторым целесообразно измерять признак e , поскольку при $e=0$ анализируется минор

$$\begin{array}{cccc|c} a & c & d & f & \\ \hline [1 & 1 & - & 0] & 1 \\ - & - & - & 1] & 2 \end{array}$$

из которого сразу видно, что признак f должен иметь значение 0, а при $e=1$ рассматривается минор

$$\begin{array}{cccc|c} a & c & d & f & \\ \hline [1 & 1 & - & 0] & 1 \\ 0 & - & 0 & - & 3 \end{array}$$

и далее измеряется признак a . При $a=0$ в миноре не останется строк с определенными значениями признака f , значит, значение последнего не может быть определено. А при $a=1$ минор сокращается до единственной строки

$$\begin{array}{ccc|c} c & d & f & \\ \hline [1 & - & 0] & 1 \end{array}$$

из которой явствует, что при $c=0$ значение признака f останется неопределенным, а при $c=1$ этот признак получит значение 1.

Напомним еще раз, что при построении алгоритма распознавания мы полагаемся на справедливость имплекативных закономерностей, заданных матрицей T , считая, что новый наблюдаемый объект должен безусловно подчиняться им — в противном случае потребова-

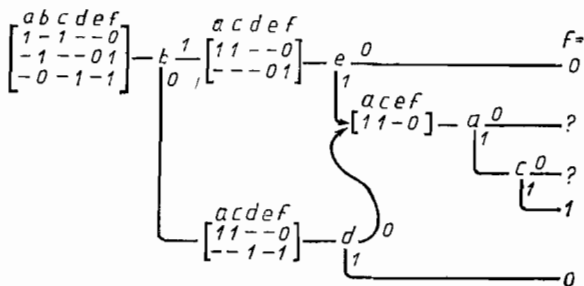


Рис. 5.11

лась бы специальная проверка этого условия, как делалось ранее. Само построение алгоритма распознавания можно при этом (когда закономерности не нарушаются) значительно упростить, с самого начала выбросив из рассмотрения все такие строки в матрице T , в которых не показаны конкретные значения целевого признака.

Для нашего примера это будет означать, что анализу подвергается минор

$$\begin{array}{c} a \quad b \quad c \quad d \quad e \quad f \\ \left[\begin{array}{cccccc} 1 & - & 1 & - & - & 0 \\ - & 1 & - & - & 0 & 1 \\ - & 0 & - & 1 & - & 1 \end{array} \right] \end{array}$$

Процесс построения алгоритма распознавания для данного примера отображен на рис. 5.11, где показано «не совсем дерево» — две ветви, оказавшиеся одинаковыми, для компактности рисунка заменены одной.

Правда, такое упрощение может привести к тому, что построенный алгоритм окажется сложнее, чем он мог бы быть...

Мы нарочно остановились на столь подробном рассмотрении как самого процесса распознавания, так и

способов его планирования, т. е. построения алгоритма распознавания, чтобы показать возможность их полной формализации и, следовательно, автоматизации. При таком рассмотрении становится очевидным, что эти процессы может с успехом выполнять машина, и ясно, как ее можно запрограммировать для этого.

При программировании процессов распознавания можно использовать любой из следующих двух подходов, напоминающих интерпретацию и компиляцию в системном программировании.

При первом из них программа динамического распознавания моделирует приведенные выше рассуждения практически в полном их объеме. Исходными данными

для нее служат матрица импликативных закономерностей T и указание на целевой признак. При работе программы выполняется последовательность операций выбора очередных признаков для измерения и анализа соответствующих миноров матрицы T , завершаемая получением окончательного результата: выдачей значения целевого признака или отказом от распознавания. Такая программа-интерпретатор универсальна — она годится для любой матрицы T и любого целевого признака, которые можно задавать перед новым выполнением программы.

При втором подходе процесс распознавания разбивается на два этапа. Сначала по матрице T и указанию на целевой признак строится некоторая конкретная программа распознавания, настроенная на эти данные, а потом уже она может выполняться, распознавая новые объекты того же класса. Такая программа реализует, к примеру, определение значения признака f по правилам, наглядно представленным на рис. 5.12. Она оказывается значительно проще, да и работает гораздо быстрее, но достигается это ценой потери универсальности — программа годится лишь для заданных конкретных матрицы T и целевого признака. При смене этих данных потребуется вернуться к первому этапу — построению новой конкретной программы распознавания, а такое построение сложнее, чем описанная выше интерпретация.

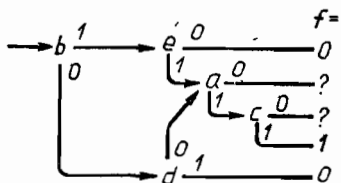
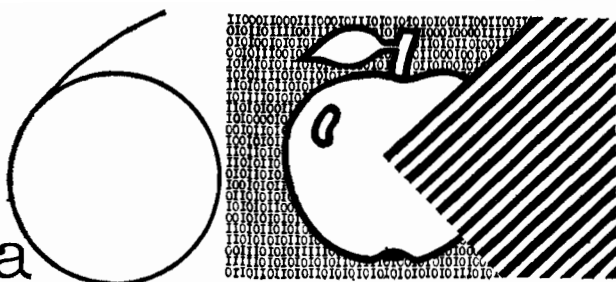


Рис. 5.12

Первый подход, интерпретационный, целесообразен в неустойчивой ситуации, когда достаточно часто меняется установка на целевой признак, а иногда и матрица имплицативных закономерностей. Тогда при каждой конкретной паре значений этих величин распознаванию подлежит небольшое число объектов, в пределе — один. Лишь при значительном росте этого числа более эффективным окажется второй подход, компилятивный. При его применении дополнительные расходы машинного времени на конструирование специализированной программы распознавания, настроенной на конкретную ситуацию, окупаются существенным ускорением процедуры распознавания рассматриваемых затем объектов.

Короче говоря, лучшим оказывается тот из подходов, который быстрее приводит к цели.



КОГДА ИЗМЕРЕНИЯ ЧАСТИЧНЫ

Рассмотренные выше процессы распознавания в сильной степени идеализованы. Обычно распознавание ведется в условиях получения лишь частичной, ограниченной информации об анализируемых объектах. К примеру, такая ситуация характерна для любых живых организмов, познавательные возможности которых ограничены определенным уровнем сложности нервной системы и наличием тех или иных анализаторов. Между тем ориентироваться приходится в значительно более сложном мире, чем то его отражение, которое синтезируется органами чувств и нервной системой организма.

Предположим, однако, что множество признаков S позволяет адекватно описывать объекты исследуемого класса. Но даже в этом случае может оказаться, что на этапе обучения у входящих в обучающую выборку объектов измеряются (или могут быть измерены) не все признаки, а лишь некоторые, иногда немногие. Иначе говоря, представляющая такую выборку матрица оказывается уже не булевой, а трюичной: кроме элементов 0 и 1 в ней фигурируют символы «—», отмечающие признаки, значение которых у соответствующих объектов остается неизвестным.

Ранее подобные матрицы использовались для перечисления пустых интервалов булева пространства признаков. Теперь матрица интерпретируется иначе: она представляет собой перечень интервалов, относительно которых известно, что они не пусты.

Например, такая матрица может начинаться со следующих строк:

$$T = \begin{array}{c} \begin{array}{cccccc} a & b & c & d & e & f \end{array} \\ \begin{bmatrix} 1 & 0 & - & 1 & - & - \\ - & 0 & 1 & - & 1 & 1 \\ 0 & - & - & 1 & 0 & - \\ - & - & 0 & - & 0 & 1 \end{bmatrix} \end{array} \begin{array}{l} 1 \\ 2 \\ 3 \\ 4 \\ \dots \end{array}$$

Первая строка означает, что в выборку попал некий объект с признаками a и d , но без признака b . Обладает ли этот объект остальными признаками c , e , f , неизвестно. Другими словами, при наблюдениях за объектом удалось установить лишь то, что он находится в интервале пространства признаков, представляемом троичным вектором

$$1 \ 0 \ - \ 1 \ - \ -$$

Следовательно, этот интервал не пуст. Аналогичная информация содержится в остальных строках. Каждая из них представляет собой вытекающее из результатов эксперимента утверждение о непустоте некоторого интервала.

Как обрабатывать подобные матрицы? Как извлекать из них закономерности?

Гипотезы об имплекативных закономерностях выдвигаются, когда в пространстве признаков обнаруживаются пустые интервалы относительно небольшого ранга. Теперь эту информацию следует получить из системы утверждений о непустоте некоторых интервалов.

Прежде всего уточним понятие пустого интервала. Здесь возможны две крайние точки зрения.

При одной из них считается, что некоторый интервал пуст, если он не пересекается ни с одним из интервалов, заданных строками матрицы эксперимента. В самом деле, какие бы значения ни принимали неизмеренные признаки у вошедших в обучающую выборку объектов, последние никак не могут попасть в рассматриваемый интервал.

Эта точка зрения слишком жесткая и потому малопродуктивна. Достаточно допустить существование хотя бы одного объекта, у которого по какой-то причине не измерен ни один признак, чтобы в соответствии с этим подходом признать, что пустых интервалов нет.

Следовательно, нет и оснований для выдвижения гипотез о каких-то закономерностях. И уже четырех показанных в матрице T строк достаточно, чтобы вычеркнуть из списка претендентов на пустые все интервалы первого и второго ранга, кроме двух:

$$\begin{array}{cccccc} - & - & - & 0 & - & 0, \\ - & 1 & - & - & 1 & - . \end{array}$$

Другая точка зрения, напротив, максимально мягкая, гибкая. При ней интервал считается пустым, если не известно с полной достоверностью, что в него попал по меньшей мере один из объектов обучающей выборки. Значит, интервал пуст, если он не включает в себя ни один из интервалов, заданных строками матрицы эксперимента.

Данное определение связано с допущением: неизмеренные признаки у объектов обучающей выборки могут принять такие значения, что все эти объекты окажутся за пределами рассматриваемого интервала. Согласно этой точке зрения, четыре первые строки в матрице T из 12 интервалов первого ранга отбраковывают только 9. Три интервала

$$\begin{array}{cccccc} a & b & c & d & e & f \\ \left[\begin{array}{cccccc} - & 1 & - & - & - & - \\ - & - & - & 0 & - & - \\ - & - & - & - & - & 0 \end{array} \right] \end{array}$$

признаются пока (до рассмотрения последующих объектов обучающей выборки) пустыми. А из интервалов второго ранга, общим числом 60 (как следует из формулы $2^2 \cdot c_n^2$, где n — число признаков), отбраковывается лишь 15. Остается 45.

Например, первая строка (1 0 — 1 — —) матрицы T означает непустоту интервалов второго ранга

$$\begin{array}{cccccc} a & b & c & d & e & f \\ \left[\begin{array}{cccccc} 1 & 0 & - & - & - & - \\ 1 & - & - & 1 & - & - \\ - & 0 & - & 1 & - & - \end{array} \right] \end{array}$$

Так в пространстве признаков определяются все пустые интервалы ограниченного ранга. Скажем, при 6 признаках и 50 объектах в обучающей выборке нет смысла рассматривать интервалы с рангом более 2. Затем оценка достоверности соответствующих импликативных гипотез дифференцируется, уточняясь для каждой гипотезы. С этой целью подсчитывается число элементарных экспериментов, подтверждающих гипотезу. Если же какой-либо эксперимент противоречит гипотезе, то последняя отвергается (и в этом случае отрицание более значимо, чем подтверждение).

Элементарный эксперимент заключается в выборке одного объекта на этапе обучения. Естественно считать, что он подтверждает гипотезу, если соответствующие интервалы не пересекаются.

Допустим, мы оцениваем следующие гипотезы, заданные в форме интервалов:

$$\begin{array}{cccccc} a & b & c & d & e & f \\ \left[\begin{array}{cccccc} 0 & 0 & - & - & - & - \\ 0 & 1 & - & - & - & - \\ 1 & 0 & - & - & - & - \\ 1 & 1 & - & - & - & - \end{array} \right] & \begin{array}{l} 1 \\ 2 \\ 3 \\ 4 \end{array} \end{array}$$

сравнивая их со строками матрицы эксперимента

$$T = \begin{array}{cccccc} a & b & c & d & e & f \\ \left[\begin{array}{cccccc} 1 & 0 & - & 1 & - & - \\ - & 0 & 1 & - & 1 & 1 \\ 0 & - & - & 1 & 0 & - \\ - & - & 0 & - & 0 & 1 \end{array} \right] & \begin{array}{l} 1 \\ 2 \\ 3 \\ 4 \end{array} \end{array}$$

...

Оказывается, что на этом начальном участке матрицы гипотеза 1 подтверждается лишь первым элементарным экспериментом. Соответствующие им интервалы не пересекаются, поскольку представляющие их троичные векторы ортогональны: один из них имеет в компоненте a значение 0, а другой — значение 1. В то же время гипотеза 2 подтверждается двумя экспериментами, а гипотеза 4 — тремя. Гипотеза 3 отвергается, поскольку ей противоречит эксперимент 1: при нем устанавливается существование объекта, принадлежащего интервалу $1\ 0\ -\ -\ -\ -$.

В результате таких расчетов некоторые гипотезы отвергаются, а для каждой из остальных определяется число подтверждающих ее элементарных экспериментов, которое естественно оказывается меньше общего числа экспериментов, что приводит к некоторому понижению достоверности гипотез. Если это понижение будет значительным, от соответствующих гипотез придется отказаться.

Описанный метод вполне применим при решении другой задачи, известной под названием «заполнение пропусков в экспериментальных таблицах». Исходная информация здесь получается совершенно аналогично. Из исследуемого класса выбираются случайно некоторые объекты, у них измеряются значения отдельных признаков. Результаты измерений оформляются в виде троичной матрицы эксперимента T . Ряд элементов матрицы получает значение «—», соответствующее неизмеренным признакам. Они и называются пропусками в таблице. Требуется их заполнить, догадавшись, какие действительные значения скрываются под символами неопределенности.

Ситуация схожа с той, в которой оказывается читатель старой книги, истрепанной настолько, что в ней можно разобрать далеко не все буквы и слова.

Задача о заполнении пропусков в экспериментальных таблицах сводится к распознаванию признаков и может быть решена теми же средствами: сначала находятся все имплекативные закономерности с достаточно высокой достоверностью, а затем они используются для вычисления недостающих значений в матрице. При этом некоторые элементы могут так и остаться неопределенными — если среди полученных закономерностей не отыщется такая, по которой можно было бы вычислить их значения.

Например, пусть в результате анализа матрицы эксперимента

$$T = \begin{array}{c} \begin{array}{cccccc} a & b & c & d & e & f \end{array} \\ \left[\begin{array}{cccccc} 1 & 0 & - & 1 & - & - \\ - & 0 & 1 & - & 1 & 1 \\ 0 & - & - & 1 & 0 & - \\ - & - & 0 & - & 0 & 1 \end{array} \right] \\ \dots \end{array}$$

найлены три имплекативные закономерности, представленные следующими векторами:

a	b	c	d	e	f
0	—	0	—	—	—
—	—	—	1	1	—
—	0	—	—	—	0

(допустим, что эти закономерности подтверждаются достаточно многими из непоказанных здесь строк матрицы T и не опровергаются ни одной из них).

Тогда матрица T доопределяется до следующего вида:

$$T = \begin{bmatrix} a & b & c & d & e & f \\ 1 & 0 & — & 1 & 0 & 1 \\ — & 0 & 1 & 0 & 1 & 1 \\ 0 & — & 1 & 1 & 0 & — \\ 1 & — & 0 & — & 0 & 1 \\ \dots & & & & & \end{bmatrix}$$

Однако после такого дополнения матрица T уже не может играть роль чисто экспериментальной. Поиск каких-либо новых гипотез по ней был бы ошибкой. С некоторой натяжкой можно сказать, что при этом некоторые исходные данные были бы учтены дважды: в оригинале, в виде первичных материалов, и в виде следствий.

В рассмотренном методе, как и ранее, решение задачи распадается на два этапа: предварительный — индуктивный и последующий — дедуктивный. Сначала по данным эксперимента строится некая общая картина — система закономерностей, которая и образует основу дедуктивной системы логического вывода, работающей на втором этапе. Однако в этом случае дедуктивная система может оказаться не такой уж безупречной.

Извлечение гипотез из троичной матрицы эксперимента может сопровождаться интересным эффектом: выдвигаемые гипотезы могут обладать достаточно высокой достоверностью каждая сама по себе, но взятые вместе — входить в противоречие.

Происходит это, в частности, тогда, когда область запрета, образуемая совокупностью имплекативных закономерностей, заведомо содержит некоторый объект обучающей выборки. Пусть известно, что последний находится в интервале I пространства признаков, а законо-

мерностям соответствуют запретные интервалы $I_1, I_2, \dots, I_k$. Противоречие возникает, если выполняется отношение

$$I \subseteq I_1 \cup I_2 \cup \dots \cup I_k.$$

Пусть, к примеру, в обучающую выборку входит объект, у которого измерены признаки a, c и d , и при этом оказалось, что $a=1, c=0$ и $d=1$. Тогда будут противоречивы две закономерности $a\bar{b}$ (не существует объектов с признаком a и в то же время без признака b) и $b\bar{c}$.

Действительно, нетрудно видеть, что объединение интервалов

a	b	c	d	e
1	0	—	—	—
—	1	0	—	—

содержит интервал

$$1 \text{ — } 0 \text{ } 1 \text{ — }.$$

Другими словами, каково бы ни было значение неизмеренных признаков у рассматриваемого объекта, он попадает или в один интервал запрета, или в другой.

Может оказаться, что никакие две закономерности из рассматриваемой группы не будут вступать в противоречие, но бóльшая их совокупность будет противоречивой.

А что делать в том случае, когда множество признаков оказывается недостаточным для того, чтобы можно было ограничиться чисто логическими методами выявления закономерностей и распознавания? Такая ситуация может возникнуть при следующих обстоятельствах.

Допустим, исследуемый класс объектов с исчерпывающей полнотой описывается в пространстве признаков S^* , которые взяты в достаточном количестве, чтобы каждый объект был описан своей, присущей только ему комбинацией признаков. Однако лишь некоторые признаки удобны для измерений на этапе обучения. Они-то и образуют множество $S : S \subseteq S^*$.

Число таких признаков может быть относительно невелико. Поэтому в соответствующей экспериментальной матрице некоторые строки могут обладать одинаковыми значениями.

В таком случае целесообразно «приводить подобные», подсчитывая число строк с одинаковым значением и

оформляя результаты в виде матрицы с различными строками и дополнительным столбцом. В нем будет показано, сколько раз встречается в исходной матрице то или иное значение.

Может случиться и так, что в результаты проводимых на этапе обучения измерений вкрадутся ошибки, своеобразный «шум», приводящий к тому, что некоторые закономерности окажутся замаскированными: соответствующие им интервалы пространства признаков будут не пустыми, а лишь слабо заполненными.

В этих случаях логические методы распознавания следует дополнять статистическими, учитывающими не только наличие или отсутствие объектов с заданными наборами значений некоторых признаков, но также и их количество.

Другие постановки задачи распознавания возникают и в случае многозначных признаков, таких, например, как цвет, который может быть красным, синим, зеленым и т. д., но принимать только одно из этих значений. В этом случае булево пространство признаков заменяется пространством многозначных переменных и вместо булевых функций используются конечные предикаты. Основные понятия теории булевых функций (элементарная конъюнкция, интервал пространства, дизъюнктивная нормальная форма и др.) при этом обобщаются, равно как и методы эквивалентных преобразований формул и их минимизации.

Пусть, например, объекты исследуемого класса описываются с помощью трех переменных a , b и c , причем первая из них принимает значения из множества $\{a_1, a_2\}$, вторая — из $\{b_1, b_2, b_3, b_4\}$, третья — из $\{c_1, c_2, c_3\}$. Конкретные объекты представляются соответствующими комбинациями значений, образующими в совокупности пространство с 24 элементами ($2 \cdot 4 \cdot 3 = 24$).

Как элементы, так и интервалы этого пространства удобно отображать секционированными булевыми векторами. Например, вектор 01.0010.100 задает объект, для которого $a=a_2$, $b=b_3$ и $c=c_1$, а вектор 11.1010.110 — интервал, содержащий 8 элементов, в которых a принимает значение a_1 или a_2 , $b=b_1$ или b_3 , $c=c_1$ или c_2 . Данный вектор может интерпретироваться и как элементарная конъюнкция, которая принимает значение 1 на элементах заданного интервала и только на них.

Представляя отдельные элементарные запреты секционированными булевыми векторами, можно отобразить всю область запрета секционированной булевой матрицей. Задачи распознавания сводятся к выделению в ней некоторых миноров и проверке их на вырожденность.

Впрочем, эти вопросы выходят за пределы настоящей книги, ведь в ней рассматривается подход к чисто логическому распознаванию в пространстве двоичных переменных. Однако основные проблемы распознавания отражаются и в рамках этого подхода, который может послужить хорошим логическим основанием для развития практических методов распознавания, реализуемых в автоматических системах с «разумным» поведением.

ОГЛАВЛЕНИЕ

Введение	3
ГЛАВА 1	
Надежное распознавание — основа разумных действий	5
ГЛАВА 2	
Связи в пространстве признаков	17
ГЛАВА 3	
Закономерности и их выявление	46
ГЛАВА 4	
Силлогистическая модель мира	67
ГЛАВА 5	
Распознавание как логический вывод	87
ГЛАВА 6	
Когда измерения частичны	108

Научно-популярное издание

ЗАКРЕВСКИЙ АРКАДИЙ ДМИТРИЕВИЧ

ЛОГИКА РАСПОЗНАВАНИЯ

Заведующая редакцией *Н. Т. Ломако*

Редактор *С. В. Машканова*

Художник *Д. А. Милованов*

Художественный редактор *А. А. Кулаженко*

Технический редактор *С. А. Курган*

Корректор *А. А. Баранова*

ИБ № 3290

Печатается по постановлению РИСО АН БССР.

Сдано в набор 01.02.88. Подписано в печать 28.04.88.
АТ 13593. Формат 84×108¹/₃₂. Бум. тип. № 2. Гарнитура
литературная. Высокая печать. Усл. печ. л. 6,30. Усл.
кр.-отт. 6,62. Уч.-изд. л. 5,81. Тираж 6320 экз. Зак. № 304
Цена 25 к.

Издательство «Наука и техника» Академии наук БССР
и Государственного комитета БССР по делам издательств,
полиграфии и книжной торговли. 220600. Минск, Ленин-
ский проспект, 68. Типография им. Франциска Скорны
издательства «Наука и техника». 220600. Минск, Ленин-
ский проспект, 68.