

УДК 519.95:681.3.06

ЛОГИЧЕСКИЕ ФУНКЦИИ В ЗАДАЧАХ ЭМПИРИЧЕСКОГО ПРЕДСКАЗАНИЯ

Г.С. Лбов

В данной работе рассматриваются алгоритмы решения различных задач, возникающих при обработке эмпирических таблиц. Под эмпирической таблицей понимается таблица, элементы которой есть результаты замеров ряда признаков у некоторого подмножества объектов, выбранных из генеральной совокупности (универсума) Γ . Любому объекту $a \in \Gamma$ можно сопоставить вектор $x = (x_1, \dots, x_j, \dots, x_n)$ и значение x_0 в пространстве признаков $X_1, \dots, X_j, \dots, X_n$; X_0 . Признак X_0 выделен в качестве целевого признака. Для каждого признака определена область его значений D_j ($j=1, \dots, n$; 0) и указана тип шкалы, в которой он замерен^{*)} [1].

Пусть выбрано некоторое множество объектов $A = \{a_1, \dots, a_i, \dots, a_N\}$, $A \subset \Gamma$. Множеству A соответствует эмпирическая таблица $v = \{x_{ij}\}$ ($i = 1, \dots, N$, $j = 1, \dots, n$, 0). Таблица v используется для решения следующих типов задач эмпирического предсказания.

1. Распознавание образов (предсказание значения целевого признака x_0 для любого объекта $a \in \Gamma$ по его описанию x ; в этом случае признак X_0 замерен в шкале наименований).

2. Упорядочивание объектов по их перспективности с точки зрения некоторого критерия (предсказание порядка на объектах некоторого подмножества A' ($A' \neq A, A' \subset \Gamma$)).

3. Предсказание значения целевого признака x_0 для $a \in \Gamma$ по его описанию x . Признак X_0 — порядковый или количественный. Заметим, что если $X_1, \dots, X_n, X_0$ — количественные признаки, то рассматриваемая задача совпадает с классической задачей восстановления функции $x_0 = f(x)$.

*) Для дальнейшего важно будет различать группы шкал, предназначенных для измерения следующих трех видов признаков: количественных (шкалы интервалов, отношений и абсолютная); порядковых (шкалы порядка, частичного порядка, рангов, баллов); номинальных (шкалы наименований).

4. Автоматическая группировка объектов. В данном случае значения признака X_0 для подмножества объектов A не заданы. Необходимо эти значения определить, используя свойство "похожести" объектов.

5. Динамическое прогнозирование значения целевого признака X_0 , использующее временные изменения в значениях признаков $X_1, \dots, X_n$.

6. Адаптивное планирование экспериментов, и их обработка при поиске приближенного значения глобального экстремума функции $x_0 = f(x)$ (на каждом шаге адаптации на основе эмпирической таблицы осуществляется прогноз функции предпочтения, задаваемой в области поиска экстремума).

Основное внимание в данной работе уделяется алгоритмам решения указанных типов задач в случае эмпирических таблиц, характеризующихся либо большим числом признаков ($n=30-150$) и малым объемом выборки ($N=n$), либо разнотипностью признаков $X_1, \dots, X_n$ (в том и другом случае в таблицах могут быть пропуски). Таблицы, для которых выполняется хотя бы одно из этих свойств, будем называть РП-таблицами. Такого рода таблицы часто встречаются при комплексных исследованиях в области медицины, геологии, социологии. Известные методы решения задачи предсказания ориентированы в основном на тот случай, когда признаки $X_1, \dots, X_n$ либо булевы, либо количественные. Применение для разнотипных признаков известных алгоритмов, использующих гипотезу "близости", "компактности" и т.д., связано с методологическими трудностями: при вычислении подобия (расстояния) между двумя векторами приходится оперировать с его компонентами, которые являются результатами измерения очевидно несравнимых величин. Поэтому в данной работе в случае разнотипных признаков для решения указанных задач предлагается использовать класс логических и линейно-логических решающих правил [2,4]. Кроме того, из теоретических исследований [3-5] следует, что данный класс правил обладает наименьшей мерой сложности по сравнению с известными классами (например, по сравнению с классом потенциальных функций, классом полиномов), что позволяет строить надежные решающие правила при большой размерности пространства признаков и малой обучающей выборке (малом числе объектов). Это обстоятельство приводит к необходимости использовать указанный класс правил также в том случае, когда признаки $X_1, \dots, X_n$ измерены в одной шкале.

Необходимо отметить, что постановка и решение задач автоматической группировки объектов, динамического прогнозирования и поиска экстремума для случая развитых признаков формулируется впервые.

При решении первых пяти типов задач используется следующая основная эмпирическая гипотеза: считается, что подмножество объектов A из множества G выбирается случайным образом. Все известные методы предсказания строят решающее правило, обладающее максимальным качеством прогноза на объектах подмножества A . Но только в случае принятия данной гипотезы можно показать, что такое решающее правило будет обладать наилучшей прогнозирующей способностью на остальных объектах множества G .

При решении задачи шестого типа использовалась адаптивная стратегия планирования экспериментов, предложенная в работе [6] для поиска приближенного значения глобального экстремума функции.

§1. Обзор различных подходов к решению задачи распознавания образов

Рассмотрим существующие подходы при решении задач распознавания образов. Одно из направлений работ связано с предварительным восстановлением в пространстве признаков функций распределения вероятностей. При первом подходе вид этих функций задается из некоторых априорных соображений. Параметры распределения оцениваются по выборке. Решающее правило определяется из отношения правдоподобия. При втором подходе вид функций распределения вероятностей не задается. Вводится лишь ограничение на "вариабельность" распределения вероятностей в виде некоторого параметра сглаживания. Этот подход объединяет непараметрические методы (метод "парзеновского окна", метод "ближайших соседей", метод "потенциальных функций" [7] и т.д.), обзор которых дается в работе [8]. В случае булевых признаков (третий подход) для аппроксимации распределения вероятностей в качестве базисных функций используются полиномы Радемахера-Уолша [8].

При другом направлении работ ограничение на класс распределений вероятностей в явном виде не задается. Решающее правило определяется из некоторого параметрического класса. Вектор параметров подбирается таким, чтобы критерий качества распознавания принимал экстремальное значение. В зависимости от выбранного класса решающих правил, критерия качества распознавания и метода поиска

экстремума имеется большая разновидность алгоритмов. Первый подход этого направления объединяет алгоритмы выбора решающих правил из класса простых функций при количественных признаках. В этом случае наиболее подробно исследован класс линейных [8, гл.5; 3, ч.3] и кусочно-линейных решающих правил [9,10] и др. В работе [11] описан алгоритм построения правила в виде набора гиперсфер (алгоритм ДРЭТ).

В рамках второго подхода были предложены методы построения решающих правил для признаков, замеренных в разных шкалах. В работе [12] вводится понятие "меры сходства" между реализациями с учетом типа признака. Рассматриваются три типа признаков: булевы, номинальные и количественные. "Мера сходства" в признаковом пространстве вводится как сумма "мер сходства" по каждому из признаков, входящих в сумму со своими "весовыми коэффициентами". Задача заключается в том, чтобы подобрать те значения параметров ("весовых коэффициентов"), при которых достигается минимальный разброс внутри каждого из образов и максимальное "расхождение" между образами.

В работе [13] были предложены алгоритмы распознавания, основанные на вычислении оценок. В этом случае также вводится некоторый вектор параметров $(\epsilon_1, \dots, \epsilon_n, \delta)$, где ϵ_i - порог "неразличимости" реализаций. Если для распознаваемой реализации при любой фиксированной подсистеме признаков число "неразличимых" реализаций рассматриваемого образа превышает порог δ , то этой реализации приписывается один "голос" в пользу этого образа. Значения параметров подбираются такими, при которых некоторый критерий, связанный с числом ошибок распознавания, принимает минимальное значение. При принятии решения используется процедура голосования по различным подмножествам признаков.

Третий подход сводится к построению по булевым признакам решающих функций алгебры логики (например, алгоритм "Кора" [14], метод "тупиковых тестов" [15]). Наиболее полный обзор методов данного подхода приводится в работе [16].

Рассмотрим основные трудности, возникающие при построении решающего правила на основе РП-таблиц.

Первый подход, связанный с восстановлением функций распределения вероятностей, предполагает, что вид этих функций известен. Но предположение о каком-либо законе распределения (например, нормальном) всегда связано с вопросом об адекватности предполагаемо-

го закона действительному. Во многих практических задачах приходится иметь дело с многомерными плотностями распределения. Однако при малом объеме выборки и большой размерности пространства невозможно надежно восстановить неизвестную функцию распределения. Например, метод оценки плотностей распределения с помощью "парзеновского окна" при числе признаков $n = 10$ для надежного построения решающего правила требует объем выборки $N > 400$ [17, с. 88]. Объем выборки, необходимый для надежного обучения, растет экспоненциально с увеличением размерности пространства. При практическом применении известного метода потенциальных функций встречаются те же трудности [8, с.194].

При большой размерности аналогичные проблемы возникают и в рамках второго направления работ (в случае поиска наилучшего правила из некоторого фиксированного класса решающих правил). Из теоретических исследований [3-5] следует: чем класс сложнее, тем больше вероятность (при минимизации числа ошибок на обучающей выборке, найти плохое решающее правило в смысле вероятности ошибки распознавания новых объектов.

Проблема построения решающего правила при большой размерности и малой выборке сводится к выбору класса правил, обладающего малой мерой сложности. В то же время класс должен быть достаточно богатым, чтобы можно было эффективно решать прикладные задачи. Малой мерой сложности обладает класс функций алгебры логики первого порядка [3, с. 117]. Этот класс функций был использован в алгоритме "Кора". Однако предварительное сведение разнотипных признаков к булевым, как правило, связано с потерей информации, иногда весьма существенной.

Поэтому в работах [2,4] был предложен класс логических и линейно-логических решающих правил (для признаков, замеренных в разных шкалах), обладающий малой мерой сложности. Даже если все признаки исходной системы являются количественными, в условиях малой выборки и большой размерности пространства использование данного класса является предпочтительным по сравнению с известными классами из-за его малой меры сложности. Кроме того, этот класс решающих правил позволяет строить простые оптимизационные процедуры поиска наилучшего решающего правила в отличие от метода, использующего "меру сходства" [12], и метода, основанного на вычислении оценок [13] (подбор параметров решающего правила в этих методах приводит к решению сложных многоэкстремальных задач).

§2. Требования к классу решающих правил

Известно, что качество решающего правила распознавания зависит от информативности признаков исходной системы, от априорных предположений о виде распределений или выбора класса решающих правил, от объема и репрезентативности выборки, от выбора критерия качества и от процедуры оптимизации этого критерия.

Сформулируем основные требования к классу решающих правил, определяемых на основе ПТ-таблиц.

Класс правил должен:

- 1) иметь малую меру сложности, но в то же время быть достаточно богатым для эффективного решения прикладных задач;
- 2) быть инвариантным к допустимым преобразованиям явля;
- 3) содержать легко интерпретируемые решающие правила;
- 4) позволять строить простые оптимизационные процедуры поиска наилучшего правила;
- 5) содержать технически просто реализуемые правила;
- 6) позволять строить алгоритмы, работающие при наличии пропусков в эмпирических таблицах.

Остановимся подробнее на каждом из этих требований.

В работе [18] впервые был рассмотрен вопрос о связи между сложностью и надежностью распознавания. В частности, была показана возможность ухудшения качества правила распознавания, построенного на основе малой выборки, с увеличением размерности пространства признаков на примере нормального распределения с единичными матрицами ковариации.

Напомним статистический критерий, который использовался при решении этого вопроса. Пусть заданы распределения $\{P(\omega, x) = P(\omega) P_{\omega}(x)\}$ на множестве признаков (ω - номер образа, x - точка n -мерного пространства признаков, $P(\omega)$ - априорная вероятность образа ω ; $\omega = 1, \dots, k$, k - число образов). Можно говорить, что набор вероятностей $s = \{p(\omega, x)\}$ задает стратегию природы ($s \in S$, S - множество стратегий природы). Обозначим через V множество всевозможных выборок объема N . Конкретная эмпирическая таблица представляет собой элемент этого множества ($v \in V$). Таблице v ставится в соответствие правило распознавания f из некоторого класса Φ . Это правило выбирается в соответствии с алгоритмом обучения Q , т.е. $Q(v) = f$. Оператор Q представляет собой либо процедуру оценки параметров распределений $\{P(\omega, x)\}$, либо оптимизационную процедуру

выбора правила f из класса Φ , минимизирующую критерий качества распознавания (например, число ошибок распознавания). Правило f дает разбиение μ области $D = D_1 \cup \dots \cup D_j \cup \dots \cup D_n$ пространства признаков на непересекающиеся области $D^1, \dots, D^j, \dots, D^k$ ($D \cap D^w = D, D^w \cap D^x = \emptyset$). Решающее правило есть отображение, сопоставляющее точке x этого пространства значение $w \in \{1, \dots, k\}$. Множеству решающих правил Φ соответствует множество разбиений Ψ пространства признаков. Для разбиения μ (фиксированного правила f) может быть определена вероятность ошибки P .

Для каждой стратегии природы s может быть задано распределение вероятностей на множестве выборок $P(v)$. Значит, решающее правило f будет выбираться случайно из множества Φ в соответствии с некоторым распределением $p(f)$. Заметим, что в общем случае $p(f) \geq p(v)$, так как алгоритм Q может выбирать одно и то же решающее правило f для некоторого подмножества выборок $V_f \subset V$. Вероятность ошибки P будет случайной величиной с некоторым распределением вероятностей $\phi(P)$. При увеличении объема выборки ($N \rightarrow \infty$) величина $P \rightarrow P_\infty$ и распределение $\phi(P)$ вырождается в δ функцию.

Для фиксированного распределения $\phi(P)$ можно определять вероятность $\text{Вер} \{P \in [P_\infty, P_\gamma]\} \geq \gamma$, ($P_\gamma \geq P_\infty$). Величина γ близка к единице. Для фиксированного значения γ из указанного неравенства можно определить значение P_γ . Величина $\epsilon = P_\gamma - P_\infty$ отражает степень отклонения выборочного правила от оптимального в смысле вероятности ошибки при фиксированной стратегии природы s , при фиксированном объеме обучающей выборки N , при фиксированном классе решающих правил. Так как стратегия природы s неизвестна, то будем рассматривать величину $\epsilon^* = \max_{s \in \Theta} \epsilon$. Величина ϵ^* при фиксированном объеме вы-

борки зависит только от выбранного класса решающих правил.

ОПРЕДЕЛЕНИЕ. Из двух классов правил Φ_1 и Φ_2 класс Φ_2 сложнее класса Φ_1 , если значение ϵ_2^* для Φ_2 больше значения ϵ_1^* для Φ_1 .

Величина ϵ^* отражает статистическую меру сложности класса Φ . Чем больше число возможных решающих правил $|\Phi|$ содержит рассматриваемый класс Φ , тем больше значение ϵ^* . Пусть, например, пространство признаков разбито на m областей. При распознавании двух образов число возможных решающих правил в этом случае равно 2^k . В работе [19] приводятся значения величины ϵ^* для различных чисел m, N, γ . При фиксированных значениях N и γ величина ϵ^* увеличивается при увеличении числа m (при увеличении числа возможных решающих правил в классе).

В работе [3] вводится понятие сложности класса решающих правил на основе равномерной сходимости частот к вероятностям

$$\text{Вер} \left\{ \sup_{f \in \Phi} |\bar{P}_f - P_f| > \varepsilon \right\} \leq 1 - \gamma,$$

где $\bar{P}_f$ - оценка вероятности ошибки на обучающей выборке для фиксированного правила f , P_f - соответствующая вероятность ошибки. Величина

$$\varepsilon = \sqrt{\frac{\ln |\Phi| - \ln (1-\gamma)}{2N}}$$

отражает статистическую меру сложности класса правил Φ . Если каждому решающему правилу f соответствует разбиение пространства признаков на m областей, то при распознавании k образов $|\Phi| = k^m |\Psi|$, где $|\Psi|$ - число различных возможных разбиений. Таким образом, чем меньше число $|\Psi|$ и число областей m , тем проще класс решающих правил. Для наиболее простого класса (класса линейных правил) $|\Phi| = 2^N$, если $N < n$, либо $|\Phi| = 2 \sum_{i=1}^{n-1} C_{N-1}^i$, если $N \geq n$ [3, с. 98].

Ниже будет показано, что рассматриваемый в данной работе класс логических решающих правил не сложнее класса линейных правил.

Остановимся на других требованиях к классу решающих правил.

Естественно следующее методологическое требование: результаты распознавания не должны зависеть от числового представления эмпирических таблиц. Пусть признакам $X_1, \dots, X_n; X_0$ соответствуют группы допустимых преобразований $E_1, \dots, E_n; E_0$. Обозначим через $\tilde{\varphi} \nu$ преобразованную таблицу ν , $\tilde{\varphi} = (\varphi_1, \dots, \varphi_n, \varphi_0)$, $\varphi_i \in E_i$. Напомним, что $Q(\nu) = f$, $f(x) = m$. Формальная запись требования инвариантности имеет вид $\varphi_0^{-1} \{ [Q(\tilde{\varphi} \nu)] (\tilde{\varphi} x) \} = [Q(\nu)](x)$.

Особое внимание необходимо обратить на требование интерпретируемости решающего правила. Полученные правила, кроме их использования для прогнозирования значений целевого признака у новых объектов, могут быть интересными для специалистов из прикладной области сами по себе. Они должны быть легко интерпретируемы. Это может оказаться полезным при формировании модели изучаемых объектов, при диалоговом режиме распознавания и т.д. Необходимо отметить, что даже простейшие решающие правила (линейные правила либо правила, задающие описание образов в виде гиперопер) не обладают этим свойством.

В том случае, когда необходимо создавать специализированное распознающее устройство, решающее правило должно быть технически просто реализуемо.

При решении медицинских, геологических, биологических и др. задач часто встречаются эмпирические таблицы с пропусками. Большинство известных алгоритмов распознавания при пропуске какого-либо значения хотя бы одного из признаков не используют соответствующую реализацию. При достаточно большом числе пропусков почти все реализации, данные для обучения, не будут использованы. Класс решающих правил должен быть таким, чтобы при построении правила можно было бы учитывать все непропущенные элементы таблицы.

§3. Класс логических решающих правил

В работе [2] был рассмотрен класс логических решающих правил для признаков, измеренных в разных шкалах. Ниже будет показано, что этот класс удовлетворяет сформулированным выше требованиям.

Введем необходимые определения и обозначения. Обозначим через $A^1, \dots, A^m, \dots, A^k$ множества объектов, данных для обучения. Множеству объектов A^m соответствует эмпирическая таблица $\{x_{ij}^m\} (i \in A^m = A)$. Под элементарным высказыванием I будем понимать значение или любое объединение значений из области D_j признака X_j , измеренного в шкале наименований, либо любой интервал значений, принадлежащий области D_j , для признака, измеренного в порядковой или количественной шкале. На основе этих высказываний формируется класс логических решающих правил. Если ввести дополнительное элементарное высказывание I вида $\bigwedge_{j=1}^n \alpha_j x_j \geq \alpha_0$ для количественных признаков

($n \leq n_1$, n_1 - число количественных признаков), то можно сформулировать класс линейно-логических решающих правил [4].

Высказыванием α будем называть конъюнкцию элементарных высказываний и их отрицаний $\alpha = I_1 \wedge \bar{I}_2 \wedge \dots \wedge I_\mu$. Под длиной высказывания понимается число входящих в него элементарных высказываний. Будем говорить, что высказывание выполняется на объекте, если каждое элементарное высказывание, входящее в α , истинно на этом объекте.

Сформулируем класс решающих правил распознавания. Для простоты будем рассматривать два образа: образ α и все остальные образы ($\bar{\alpha}$). Пусть множество объектов $A = A^{\alpha}$ и $A^{\bar{\alpha}}$ выбрано случайно из некоторого универсума $\Gamma = \Gamma^{\alpha} \cup \Gamma^{\bar{\alpha}}$.

Обозначим через $R = \{a_1, \dots, a_1, \dots, a_M\}$ набор конъюнкций, дающий разбиение множества Γ , т.е. такой их набор, при котором для любого объекта $a \in \Gamma$ выполняется одна и только одна конъюнкция из этого набора.

Простейшим примером такого набора может служить набор конъюнкций, который представлен в виде дерева [4]. Числом ветвей дерева определяется число конъюнкций M , входящих в набор R . Для каждой ветви a могут быть определены число объектов $N_{a\omega}$ образа ω и число объектов $N_{a\bar{\omega}}$ образа $\bar{\omega}$, для которых выполнялась конъюнкция a . Решающее правило, минимизирующее число ошибок на обучающей выборке, задается следующим образом: если $N_{a\omega} \geq N_{a\bar{\omega}}$, то $f(a) = \omega$, иначе $f(a) = \bar{\omega}$. Конкретный набор R характеризуется числом M и относительным числом ошибок $\bar{P}$. Обозначим через $f(R)$ решающее правило, заданное на R . Множество решающих правил, заданное на всевозможных наборах $\{R\}$ при фиксированном числе M , определяет класс логических решающих правил Φ_M . При $M_1 < M_2 < \dots < M_1 < \dots$ получим вложенные классы решающих правил $\Phi_{M_1} \subset \Phi_{M_2} \subset \dots \subset \Phi_{M_1} \subset \dots$. Увели-

чивая число M для фиксированного числа признаков, увеличиваем меру сложности класса Φ_M . Для фиксированного распределения вероятностей на признаках можно определять величины: $P(f)$ — вероятность выбора решающего правила $f \in \Phi_{M_1}$ на основе обучающей выборки; P_f — вероятность ошибки распознавания при использовании правила f ; $P^1 = EP$ — математическое ожидание вероятности ошибки распознавания для класса Φ_{M_1} . Доказано [22], что при $N \rightarrow \infty \lim_{i \rightarrow \infty} P^1 = P_0$, где P_0 — вероятность ошибки при байесовском решающем правиле, определенном на признаках $X_1, \dots, X_n$.

Необходимо указать процедуру Q выбора решающего правила из класса Φ_M , для которого относительная ошибка $\bar{P}$ была бы минимальной либо при фиксированном значении $\bar{P}$ число M было бы минимальным. Такой алгоритм построения правила рассматривается в §4.

Рассмотрим несколько иной вид логического решающего правила. Поставим в соответствие некоторому набору конъюнкций $B = \{a_1, \dots, a_1, \dots, a_v, \dots, a_v\}$ новую систему $Y_1, \dots, Y_1, \dots, Y_v$ таким образом, что если для объекта a выполняется a_i , то значение признака Y_i равно 1, иначе 0. Тогда эмпирические таблицы можно переписать: $\{x_{ij}^w\} \rightarrow \{y_{1j}^w\}$ и $\{x_{ij}^w\} \rightarrow \{y_{1j}^w\}$. Для новой системы признаков можно указать множество наборов конъюнкций $\{\bar{B}\}$, каждый из которых может

быть представлено в виде дерева. В этом случае в качестве элементарных высказываний будут выступать конъюнкции из набора B . По обучающей выборке необходимо выбрать тот набор $\bar{B}^*$, для которого оценка вероятности ошибки является минимальной при заданном числе новых конъюнкций в наборе $\bar{B}^*$. Качество $\bar{B}^*$ зависит от набора $B = \{z_1, \dots, z_t, \dots, z_v\}$. В работах [2, 19] предложены алгоритмы КОРАЛЛ и ТЭМП для определения списка логических закономерностей для разнотипных признаков, который может быть предложен в качестве такого исходного набора конъюнкций. Под закономерностью, характеризующей образ w , понимается высказывание z , для которого $\bar{P}_{zw} \geq \delta$ и $\bar{P}_{z\bar{w}} \leq \beta$ (например, $\delta = 0,6$, $\beta = 0,02$). Величины $\bar{P}_{zw}$ и $\bar{P}_{z\bar{w}}$ — относительные числа объектов, на которых выполнялось высказывание z соответственно для образа w и всех остальных образов ($\bar{w}$). Для малого фиксированного β величина δ подбирается экспериментально: начиная с $\delta = 1$, эта величина уменьшается с некоторым шагом (например, $\Delta\delta = 0,05$) до тех пор, пока не обнаружится небольшое число закономерностей (обычно это число закономерностей задается равным 3–5 на образ).

Отметим ряд важных свойств введенного класса логических решающих правил для разнотипных признаков.

1. Этот класс обладает малой мерой сложности. Определим число возможных правил, представленных в виде дерева с M ветвями

$$|\Phi| = k^M \cdot |w| \leq k^M \cdot (1 \cdot n)(1 \cdot n - 1) \dots (1 \cdot n - (M - 2)) < k^M \cdot (1 \cdot n)^{M-1},$$

где $1 = \max 1_j, 1_j$ — число элементарных высказываний на признаке x_j . При $k = 2$, $M = 10$, $1 = 10$, $n = 100$ число возможных правил $|\Phi| < 10^{30}$. Для класса линейных решающих правил при $N < n$ верно $|\Phi| = 2^2 \cdot 2^{100} < 4 \cdot 10^{30}$. Таким образом, рассматриваемый класс относится к классу простейших решающих правил. Это свойство имеет первостепенное значение в условиях малой выборки и большой размерности пространства. Но, несмотря на малую меру сложности, указанный класс, как показывает опыт решения прикладных задач (см. §10), оказывается достаточно богатым для эффективного решения этих задач.

2. Можно доказать [20], что алгоритмы построения логических правил являются инвариантными к допустимым преобразованиям правил. Эти алгоритмы для любых двух таблиц, описывающих один и тот же экспериментальный материал, но представленных в разной числовой форме, обнаруживают одни и те же закономерности, видоизмененные из-за разного представления таблиц.

3. Закономерности, представленные в виде конъюнкций значений в условных интервалах признаков, легко интерпретируемы, а решающее правило может быть технически просто реализовано в виде порогового устройства.

4. Использование класса логических правил позволяет строить простейшие оптимизационные процедуры поиска наилучшего правила. Например, алгоритм ТЭМП [19] обнаруживает все логические закономерности (с точностью до выбранных элементарных высказываний), характеризующие эмпирическую таблицу. Выигрыш во времени решения по этому алгоритму по сравнению с полным перебором зависит от конкретной эмпирической таблицы. Однако, как показывает опыт решения прикладных задач, указанный выигрыш позволяет обрабатывать РП-таблицы за приемлемое машинное время.

5. Большинство известных алгоритмов распознавания при пропуске значения хотя бы одного из признаков не используют соответствующую реализацию при построении правила. При формировании логического правила эта реализация не используется только для тех конъюнкций, которые содержат признак с пропущенным значением.

6. При выборе логических закономерностей одновременно решается задача сокращения числа признаков исходной системы. Кроме того, решающее правило позволяет осуществлять "индивидуальный" подход: для каждого распознаваемого объекта используется свое подмножество признаков. Указанные свойства являются важными при решении тех задач, при которых в процессе принятия решения необходимо минимизировать, кроме ошибки распознавания, еще и стоимость измерения признаков.

7. Класс логических решающих правил обладает важным с теоретической точки зрения асимптотическим свойством: при $N \rightarrow \infty$ и $M \rightarrow \infty$ логическое решающее правило совпадает с байесовским оптимальным решающим правилом.

§4. Алгоритмы построения логических решающих правил

Алгоритм DW [4] реализует направленную процедуру поиска набора высказываний, представленного в виде дерева. На первом шаге перебора рассматриваются все элементарные высказывания для каждого из признаков исходной системы (обозначим через O_j множество элементарных высказываний на признаке X_j). Для этого по выборке определяются различные значения этого признака. Если X_j замерен в шкале наименований, то O_j состоит из различных значений признака и все-

возможных их объединений (за исключением всего множества значений x_j тех высказываний, которые являются отрицаниями уже рассмотренных в процессе перебора). Если x_j замерен в шкале порядка, отношений, интервалов, то предварительно производится разбиение области D_j на ряд интервалов так, чтобы каждый интервал содержал Δ различных значений признака (обычно значение Δ задается равным 1, 2, 3). Множество O_j состоит из указанных интервалов и всевозможных объединений соседних интервалов (ограничения аналогичны вышеуказанным).

После указанного перебора выбирается то I_j и его отрицание $\bar{I}_j$, на котором можно построить решающее правило с минимальной оценкой вероятности ошибки. Можно указать четыре подмножества объектов $A_{I_j}^w, A_{\bar{I}_j}^w, A_{I_j}^{\bar{w}}, A_{\bar{I}_j}^{\bar{w}}$. Далее перебираются все элементарные высказывания, за исключением I_j и $\bar{I}_j$. Выбирается то I_k и его отрицание $\bar{I}_k$, на которых можно построить решающее правило, минимизирующее оценку вероятности ошибки на объединенном подмножестве $A_{I_j}^w \cup A_{\bar{I}_j}^w$, а также выбирается I_l и $\bar{I}_l$, используя которые можно минимизировать ошибку на подмножестве $A_{I_j}^w \cup A_{\bar{I}_j}^{\bar{w}}$. Из двух указанных возможных условий (вершин дерева) выбирается то, которое дает меньшую ошибку. Данный процесс построения правила продолжается аналогично для каждой ветви дерева. Процедура заканчивается либо по достижении заранее фиксированной ошибки, либо по достижении заранее фиксированного числа ветвей.

Если в дереве R в качестве элементарных высказываний используются (наряду с указанными) высказывания типа $\sum_{j=1}^n \alpha_j x_j \geq \alpha_0$, то для построения решающего правила используется алгоритм, аналогичный алгоритму DW со следующими изменениями.

Алгоритм LLRP. Пусть на некотором этапе выбора наилучшей вершины дерева выбрано некоторое высказывание на количественном признаке. Обозначим этот признак через X_{j_1} . На следующем этапе в двумерных подпространствах $\{X_{j_1}, X_{j_2}\}$ ($j = 1, \dots, n_1; j \neq j_1; n_1$ — число количественных признаков в исходной системе) строятся гиперплоскости, разделяющие соответствующую данной ветви дерева обуча-

ную выборку. Для этого обучающая выборка проецируется на векторы

$$\left\{ 1, (-1)^k \cdot \operatorname{tg} \left(\frac{\pi}{2} \cdot \left(1 - \frac{\left[\frac{k+1}{2} \right]}{F+1} \right) \right) \right\},$$

где $k = 1, \dots, 2F$; 1 - координата оси X_{j_1} ; $(-1)^k \operatorname{tg} \left(\frac{\pi}{2} \cdot \left(1 - \frac{\left[\frac{k+1}{2} \right]}{F+1} \right) \right)$ - координата оси X_{j_2} ; F - число рассматриваемых векторов, которое обычно выбирается равным 5-10. Такой набор векторов в качестве новых признаков рассматривается в работе [41].

Из определенных таким образом признаков указывается тот, который минимизирует ошибку. Пусть это будет вектор с координатами $(1, \varphi_1)$. Таким образом определяется лучшая пара признаков $\{X_{j_1}, X_{j_2}\}$ и лучшее условие (гиперплоскость в пространстве этих признаков, ортогональная лучшему вектору). Если при этом ошибка уменьшается, то переходим к рассмотрению гиперплоскостей в пространствах признаков $\{X_{j_1}, X_{j_2}, X_{j_3}\}$ ($j \neq j_1, j_2$). При этом выборка проецируется на векторы с координатами $\left\{ 1, \varphi_1, (-1)^k \operatorname{tg} \left(\frac{\pi}{2} \cdot \right. \right.$
 $\left. \left. x \left(1 - \frac{\left[\frac{k+1}{2} \right]}{F+1} \right) \right) \right\}$ (в противном случае процесс построения гиперплоскостей завершен: это будет одна из гиперплоскостей, ортогональных признаку X_{j_1}). В общем случае указанный процесс построения наилучших гиперплоскостей может закончиться выбором гиперплоскости в пространстве m признаков ($m \leq n_1$). Таким образом, будет построено решающее правило в виде дерева, в котором в качестве элементарных высказываний могут быть использованы высказывания типа $\bigwedge_{j=1}^m \alpha_j x_j \geq \alpha_0$. Заметим, что при большом числе признаков и малом объеме выборки вероятность получения плохого набора конъюнкций может быть достаточно большой. Для повышения надежности принятых решения предлагается следующая процедура. После получения первого набора, заданного в виде дерева (либо набора, дающего покрытие объектов из обучающей выборки), исключаем признаки, вошедшие в него, а на остальных признаках определяем новый набор и т.д. до получения такого набора, при котором при заданном числе конъюнкций ошибка F первый раз превышает некоторый порог $\bar{F}$.

Алгоритм ЛРП [19] реализует решающее правило на основе всех выбранных наборов. Для распознавания контрольного объекта выбираются из общего списка выделенных конъюнкций те $a^1, \dots, a^r$, которые выполнялись на этом объекте. Для образа ω определяются подмножества объектов $A_{a^1}^\omega, \dots, A_{a^r}^\omega$, на которых выполнялись $a^1, \dots, a^r$.

Определяется число объектов N^ω , входящих в объединение указанных подмножеств. Указанная процедура осуществляется для каждого из образов. Контрольный объект относим к тому образу ω , для которого N^ω максимально.

Заметим, что процедура распознавания на основе голосования по конъюнкциям (как это сделано в алгоритме "Кора") является частным случаем алгоритма ЛРП: результаты распознавания указанных процедур совпадают, если подмножества $A_{a^1}^\omega, \dots, A_{a^r}^\omega$ не пересекаются и равномошны ($\omega = 1, \dots, k$). В других случаях применение процедуры голосования является методологически необоснованным из-за зависимости элементов друг от друга, которые используются при голосовании.

§5. Метод предсказания перспективности объектов

Содержательный смысл этой задачи может быть пояснен на примере упорядочения геологических объектов [21]. Пусть дано некоторое множество объектов $A' = \{a_1', \dots, a_1, \dots, a_N'\}$. Объекту a_1 соответствует его описание $x_1 = \{x_{11}, \dots, x_{1n}\}$ в пространстве признаков $X_1, \dots, X_n$. Он принадлежит к одному из двух образов, т.е. значение целевого признака для этого объекта $x_{10} = \omega$ либо $x_{10} = \bar{\omega}$. Но принадлежность его к тому или иному образу неизвестна. Необходимо упорядочить множество объектов A' по их перспективности. Под перспективностью объекта понимается вероятность его принадлежности к одному из образов (например, к образу ω). Для любого порядка указанных объектов ($\Pi: a_{\pi_1}, \dots, a_{\pi_N}$) можно определять значения критерия

$$\varphi(\Pi) = \sum_{i=1}^N \Delta(a_i).$$

где

$$\Delta(a_i) = \begin{cases} 0, & \text{если } x_{i0} = \omega, \\ \pi(a_i), & \text{если } x_{i0} = \bar{\omega}. \end{cases}$$

$\pi(a_i)$ — номер места, на которое поставлен объект a_i в порядке Π .

Упорядочивание тем лучше, чем больше $\phi(\Pi)$. Рассмотрим содержательную трактовку этого критерия. Пусть Γ - множество геологических участков. Тот участок, в котором обнаружено месторождение (образ m), приносит некоторый доход в единицу времени, начиная с момента обнаружения. Если в нем не обнаружено месторождения, то средства, затраченные на его разработку (например, бурение), оказываются потраченными впустую. Пусть за время T можно обследовать μ участков ($\mu \ll N$). Можно показать, что чем больше значение $\phi(\Pi)$, тем больше доход за время T .

Для решения задачи упорядочивания объектов используется обучающая выборка v , представляющая собой результаты замеров признаков некоторого множества объектов $A = A^w \cup A^{\bar{w}}$. Число объектов образа w равно N_w , образа $\bar{w}$ равно $N_{\bar{w}}$ ($N_w + N_{\bar{w}} = N$).

Алгоритм упорядочивания Q строит на основе обучающей выборки решающее правило r , представляющее собой отображение

$$r: (x_1, \dots, x_1, \dots, x_N) \rightarrow (t_1, \dots, t_\mu, \dots, t_N).$$

Необходимо отметить, что из-за случайности выборки алгоритм будет порождать случайные правила (случайные порядки). Величина $\phi(\Pi)$ будет случайной величиной.

Сформулируем алгоритм упорядочивания объектов следующим образом. На основе обучающей выборки определяется логическое решающее правило. В случае использования алгоритма ЛРП для каждого объекта $a_i \in A$ определяется два числа N^{w_i} и $N^{\bar{w}_i}$. Объекты множества A упорядочиваем в соответствии с критерием $\gamma = \frac{N^{w_i} + 1}{N^{\bar{w}_i} + 1}$, отражающим

отношение правдоподобия. На первое место ставим объект с максимальным значением γ , на втором месте объект со следующим значением этой величины и т.д.

Если на основе обучающей выборки определены несколько наборов конъюнкций, каждый из которых представлен в виде дерева, то можно предложить следующий критерий γ . Пусть на объекте a_1 выполнялось некоторое высказывание a_1 из первого набора, a_2 - из второго набора и т.д. Если считать, что события a_1 и a_2 независимы, то отношение правдоподобия определяется как

$$\gamma = \frac{N_{a_1, w} + 1}{N_{a_1, \bar{w}} + 1} \cdot \frac{N_{a_2, w} + 1}{N_{a_2, \bar{w}} + 1} \dots$$

Критерий γ используется в качестве критерия упорядочивания.

В §10 будут приведены примеры применения указанного метода для упорядочивания геологических участков по их перспективности на месторождение и упорядочивания шахт по степени пожароопасности.

В работе [22] доказано для аналогичных постановок задач упорядочивания, что алгоритм, использующий отношение правдоподобия, является оптимальным (дающим минимальное значение математической функции $\Phi(P)$).

§6. Предсказание значения порядкового или количественного признака

Для предсказания значения целевого признака X_0 , измеренного в шкале порядка, интервалов, отношений, в работе [23] был предложен алгоритм ПИНЧ. Признаки исходной системы $X_1, \dots, X_j, \dots, X_n$ могут быть измерены в любых шкалах.

Рассмотрим случай, когда целевой признак X_0 является количественным. Предварительно на основе множества значений этого признака $\{x_{ip}\}$, $i = 1, \dots, N$, производится разбиение его на интервалы. Для этого упорядочиваются выборочные значения и определяются расстояния между соседними значениями. Рассматривается возможность постановки первой границы на равном расстоянии от тех соседних значений, между которыми расстояние максимально. После этого проверяется число объектов, попавших в данные два интервала. Если хо-

тя бы в одном интервале ω относительное число объектов $\frac{N_\omega}{N} < \delta$, то граница не ставится, а рассматривается возможность ее постановки между другими соседними значениями, для которых расстояние максимально. В этом переборе не рассматриваются те соседние значения, для которых расстояние равно нулю (случай всех совпадающих значений не рассматривается). Пусть получено k интервалов ($k \geq 2$). Для каждого такого интервала ω ($\omega = 1, \dots, k$) определяется набор логических закономерностей.

Под логической закономерностью, характеризующей интервал ω на X_0 , будем понимать высказывание λ , для которого

$$\frac{N_{\lambda\omega}}{N_\omega} \geq \delta \quad \text{и} \quad \frac{\sum_{i \in \Gamma_1} v_i}{N_\omega} \leq \beta,$$

где Γ_1 — множество объектов, на которых выполняется высказывание λ и значения признаков X_0 этих объектов находятся вне интервала ω , v_i — некоторый "вес" i -го объекта.

"Вес" v_i выбирается из следующих соображений: чем больше расстояние ρ_i между значением x_{i0} и ближайшей границей интервала ω , тем больше должен быть "вес" этого объекта и тем менее информативным будет высказывание α .

Пусть $v_i = v(\rho_i)$ будет линейной функцией расстояния $v_i = c \cdot \rho_i + c_0$. Параметры c и c_0 определяются при следующих ограничениях:

$$v(\rho_{\max}) = d \cdot v(\rho_{\min}),$$

$$\sum_{i \in \Gamma_2} v_i = N - N_\omega,$$

где Γ_2 - множество всех объектов, для которых значения признака X_0 находятся вне интервала ω , d -параметр, определяющий отношение между максимальным и минимальным "весами" (этот параметр задается из эвристических соображений),

$$\rho_{\min} = \min_{i \in \Gamma_2} \rho_i, \quad \rho_{\max} = \max_{i \in \Gamma_2} \rho_i.$$

Параметры c и c_0 определяются путем несложных выкладок:

$$c = \frac{1}{\mu + (\rho_{\text{ср}} - \rho_{\min})}, \quad c_0 = \frac{\mu - \rho_{\min}}{\mu + (\rho_{\text{ср}} - \rho_{\min})},$$

$$\mu = \frac{\rho_{\max} - \rho_{\min}}{d-1}, \quad \rho_{\text{ср}} = \frac{\sum_{i \in \Gamma_2} \rho_i}{N - N_\omega}.$$

Для выбора логических закономерностей, характеризующих интервал ω , применяется алгоритм КОРАЛЛ.

Если целевой признак X_0 замерен в шкале порядка, то в алгоритм вносятся следующие изменения. При разбиении X_0 рассматриваются интервалы $x'_{j0} < x_0 \leq x'_{i0}$, где x'_{j0} и x'_{i0} выбираются из множества несовпадающих упорядоченных значений признака X_0 . Условья постановки границ те же самые, что и рассмотренные выше.

При использовании весовой функция $v(\rho_i)$ вместо расстояния выборочного значения x_{i0} от границ интервала ω учитывается его порядковый номер в упорядоченном множестве значений, не попавших в интервал ω . Номера присваиваются отдельно для значений слева и справа от интервала ω , начиная от его границ. Совпадающим значениям присваиваются одинаковые номера. Очевидно, весовая функция $v(\rho_i)$ будет инвариантной для всех монотонных преобразований признака X_0 .

Необходимо отметить, что в работе [24] для решения рассматриваемой в данном параграфе задачи предложен алгоритм прогнозирования в классе линейно-логических решающих функций для разнотипных признаков.

§7. Автоматическая группировка объектов

Пусть задана эмпирическая таблица $\{x_{ij}\}$ ($i = 1, \dots, N$; $j = 1, \dots, n$; N - число объектов; n - число признаков). Необходимо разбить данное множество объектов $A = \{a_1, \dots, a_1, \dots, a_N\}$ на ряд непересекающихся подмножеств (групп) так, чтобы в каждой группе оказались объекты в некотором смысле наиболее "похожие" между собой. Такого рода задачи получили название задач таксономии (кластер анализа, обучения без учителя и т.д.), для решения которых предложены эффективные алгоритмы, например [10, 11]. Эти задачи возникают часто на ранних этапах исследования, когда интересно получить некоторые начальные сведения о внутренней природе и структуре эмпирических данных. Как правило, таксономия объектов делается либо во всем исходном пространстве, либо в некотором подпространстве признаков, хотя естественно предположить, что множество объектов A может разбиваться на ряд таких групп, для каждого из которых каждая из которых требуется свое подмножество признаков^{*)}. Заметим, что мы переходим к обычной задаче таксономии, если указанные подмножества признаков для описания всех групп совпадают.

Кроме того, известны алгоритмы таксономии предназначены для решения задач в случае признаков одного типа (обычно используют количественные признаки). Вместе с тем на практике возникают задачи таксономии объектов по признакам, измеренным в разных шкалах. На необходимость решения этой задачи указывается и в работе [8, с. 238].

В данной работе делается попытка предложить на основе использования логических функций подход к решению задач автоматической группировки для разнотипных признаков. При этом для описания каждой группы объектов может быть использовано свое подмножество признаков.

Считается, что закономерные связи между объектами и признаками отсутствуют, если исходная эмпирическая таблица получена с

*) При описании объектов с целью распознавания образов этот эффект наблюдался на всех решенных прикладных задачах.

помощью случайного механизма, т.е. если ее можно рассматривать как выборку из области D исходного пространства с равномерным законом распределения.

В этом случае для любого логического высказывания a можно определить вероятность его выполнения p_a на такой "случайной" таблице. Вероятность выполнения a на N_a объектах из общего числа N равна

$$P(N_a) = C_N^{N_a} p_a^{N_a} (1-p_a)^{N-N_a}.$$

Выбор критерия предпочтения одного логического высказывания другому при поиске закономерностей для решения указанной задачи группировки объектов основан на следующей гипотезе: чем меньше величина $P(N_a)$ для высказывания a , тем больше оснований рассматривать это высказывание в качестве закономерности.

Число реализаций N_a , на которых выполняется высказывание a на "случайной" таблице, в среднем равно $N \cdot p_a$. Для целей группировки объектов будем рассматривать только те высказывания, для которых $N_a > N \cdot p_a$.

Кроме того, нас интересуют высказывания, выполняющиеся на исходной таблице не менее чем 6 раз.

Для простоты вычислений при определении порядка предпочтения на множестве высказываний будем использовать величину

$$\gamma(a) = - \frac{(N_a - N \cdot p_a)^2}{p_a(1-p_a)},$$

которая получается при аппроксимации биномиального распределения $P(N_a)$ нормальным приближением

$$f(N_a) = \frac{1}{\sqrt{2\pi N p_a(1-p_a)}} e^{-\frac{\gamma(a)}{2}}.$$

Ясно, что чем меньше величина $P(N_a)$, тем меньше и величина $\gamma(a)$.

Рассмотрим высказывания a_1 и a_2 , выполняющиеся на данной таблице N_{a_1} и N_{a_2} раз соответственно. Будем считать, что высказывание a_1 предпочтительнее a_2 , если $\gamma(a_1) < \gamma(a_2)$.

Для обнаружения наилучших по критерию $\gamma(a)$ высказываний используется алгоритм КОРАЛЛ (либо ТЕМ1). Для этого после выбора

лучшего первого высказывания a^1 исключаются те N_{a^1} объектов из N , на которых оно выполнялось. Затем определяется лучшее a^2 , которое выполняется только на дополнительном подмножестве объектов и т.д. до полного разбиения исходного множества объектов. Закономерности $a^1, a^2, \dots$ будем называть таксонами.

Необходимо отметить, что методы автоматической группировки объектов применяются, как правило, на раннем этапе исследований. Поэтому удобная для интерпретации форма представления описания различных групп объектов в виде логических высказываний имеет первостепенное значение.

§8. Метод динамического прогнозирования

В литературе при анализе динамических объектов основное внимание уделяется случаю, когда все признаки являются количественными. Однако в ряде прикладных областей (например, медицине, социологии, экономике) признаки, характеризующие динамический объект, могут быть измерены в разных шкалах. В этом случае для описания динамических закономерностей можно использовать класс логических функций.

В зависимости от цели исследования могут возникнуть различные постановки задач прогнозирования. Рассмотрим две из них. Пусть для описания объекта выбрана исходная система признаков $\{x_1, \dots, x_n, x_0\}$. Признак x_0 может принимать два значения: $\{\omega, \bar{\omega}\}$, где ω — интересующее состояние объекта (например, ω — некоторое заболевание пациента, $\bar{\omega}$ — отсутствие этого заболевания). Необходимо на основе результатов измерений признаков, полученных в моменты времени $t^{(1)}, \dots, t^{(L)}$ (для этих моментов $x_0 = \bar{\omega}$), определить для данного объекта промежуток времени Δt от момента времени $t^{(L)}$ до того момента, когда значение признака x_0 станет равным ω . Примером подобной постановки может служить задача ранней диагностики некоторого заболевания. Для решения этой задачи организуется обучающая выборка, представляющая собой временные замеры указанных признаков у N объектов. Для i -го объекта ($i = 1, \dots, N$) значения всех признаков одновременно определяются в R временных точках $t_1^i, \dots, t_2^i, \dots, t_R^i$ через равные интервалы времени Δt . Множества временных точек для рассматриваемых объектов могут не совпадать, но для каждого объекта выполняется следующее условие: считается, что в конечный момент времени $x_0(t_R^i) = \omega$, а во все предыдущие моменты $x_0(t_1^i) = \bar{\omega}$ ($i = 1, \dots, R-1$).

Выберем в качестве начала отсчета по времени некоторый произвольный момент t_R . Измерения признаков всех объектов, полученные в моменты времени t_R^1 для первого объекта, t_R^2 для второго объекта и т.д., соотнесем с этим началом отсчета. Тогда множество моментов $\{t_{R-1}^i\}$ соотнесем с моментом t_{R-1} , множество $\{t_{R-2}^i\}$ с t_{R-2} и т.д. (причем $t_i - t_{i-1} = \Delta t$).

Обучающую выборку можно представлять в виде последовательности таблиц $c_1, \dots, c_1, \dots, c_R$, где $c_1 = \{x_{1,j}^1\}$ ($i = 1, \dots, N$; $j = 1, \dots, n, 0$; $1 = 1, \dots, R$). На основе этих таблиц определяется последовательность таблиц $\Delta c_1, \dots, \Delta c_1, \dots, \Delta c_{R-1}$, отражающих изменения значений признаков между соседними моментами времени.

Для определения прогнозируемой величины Δt будут использоваться логические закономерности, обнаруженные на указанных двух последовательностях таблиц и отражающие закономерные (статистически значимые) изменения во времени значений признаков $x_1, \dots, x_n$ по мере приближения объектов к состоянию ω . Алгоритм определения динамических закономерностей и локализации объекта по времени (определения величины Δt) описан в работе [25].

Рассмотрим другую постановку задачи динамического прогнозирования.

Пусть обучающая выборка состоит из двух последовательностей таблиц. Первая последовательность представляет собой множество замеров признаков $x_1, \dots, x_n$ по времени в моменты $t_1, \dots, t_1, \dots, t_R$ у N_ω объектов, принадлежащих образу ω . Вторая последовательность — замеры признаков в те же моменты у объектов, принадлежащих образу $\bar{\omega}$. Необходимо на основе выделенных динамических закономерностей сформулировать решающее правило распознавания, позволяющее на основе результатов измерений у нового объекта признаков в моменты $t_1, \dots, t_1$ ($1 < R$) определять наиболее вероятный класс $(\omega, \bar{\omega})$, к которому он будет принадлежать в R -й момент времени.

Приведем пример подобной постановки задачи. Пусть изучается некоторый технологический процесс изготовления деталей (объектов). Для характеристики процесса в каждый момент времени можно замерить признаки $x_1, \dots, x_n$. После периода $T = t_R - t_1$, характеризующего длительность технологического процесса, определяется годность изделия $(\omega, \bar{\omega})$. За начало отсчета t_1 принимается начало процесса. Необходимо еще до окончания процесса предсказать образ, к которому будет принадлежать объект по измерениям признаков за 1 момент времени ($1 < R$) для того, чтобы ввести своевременное управле-

щее действие. В этом случае, как и в первой постановке, моменту времени t_1 ставятся в соответствие списки логических закономерностей, отличающихся таблиц s_1^w и $s_1^{\bar{w}}$ (таблицы As_1^w и $As_1^{\bar{w}}$) друг от друга. Далее можно сформулировать решающее правило на основе диамических закономерностей, полученных из указанных списков.

§9. Адаптивное планирование экспериментов и их обработка при поиске приближенного значения глобального экстремума

Рассматривается вещественная функция $x_0 = f(x_1, \dots, x_n)$. Переменные $x_1, \dots, x_n$ могут быть замерены в разных шкалах. Если x_j - количественная переменная, то разобьем область D_j (интервал $[a_j, b_j]$) на 1_j равных подинтервалов. В качестве D_j будем рассматривать множество значений x_j , представляющих средние точки выбранных подинтервалов.

Каждому множеству D_j ($j = 1, \dots, n$) сопоставим множество натуральных чисел $(1, \dots, \beta, \dots, 1_j)$. Мощность множества D равна

$N_0 = \prod_{j=1}^n 1_j$. На этом множестве задается метрика, предложенная в работе [12] для случая переменных, замеренных в разных шкалах.

Пусть $x^* = (x_1^*, \dots, x_n^*)$ - точка, в которой достигается максимум функции $f(x)$: $f(x^*) = \max_{x \in D} f(x)$.

Введем понятие ϵ -окрестности точки глобального максимума функции. Значение $f(x)$ называется ϵ -соседним значением $f(x^*)$ ($\epsilon \geq 0$, целое; $x, x^* \in D$), если в интервале $(f(x), f(x^*))$ находится ϵ различных значений функции. Значение $f(x)$ при $f(x) = f(x^*)$ будет 0-соседним значением. Назовем ϵ -окрестностью точки x^* множество точек $x \in D$, значения функции в которых не более чем ϵ -соседние значению $f(x^*)$.

В любой произвольной точке $x \in D$ может быть определено значение функции путем эксперимента или вычисления. В ходе поиска решено провести фиксированное число T экспериментов (считается, что определение значения функции связано с большими затратами, поэтому число T обычно велико). Правило восстановления T точек $x^1, \dots, x^T, \dots, x^T$ в области D назовем стратегией поиска. Необходимо выбрать такую стратегию планирования T различных экспериментов, чтобы получить значение функции, являющееся ϵ -соседним к $f(x^*)$. При этом значение ϵ должно быть минимальным. Достигнутое при по-

иске значение ϵ зависит от выбранной стратегии планирования экспериментов и меры сложности функции. Введение этой меры для упорядочивания по сложности многоэкстремальных функций является нерешенной задачей. В настоящий момент, как правило, в качестве такой меры используется константа Липшица, отражающая степень варируемости функции. Однако можно привести примеры, когда функция является "простой" для поиска глобального экстремума, но характеризуется большой константой Липшица.

В работе [6] впервые была предложена следующая адаптивная стратегия планирования экспериментов для поиска приближенного значения глобального экстремума. Множество из T экспериментов разбивается на R групп $T = r^{(1)} + \dots + r^{(\varphi)} + \dots + r^{(R)}$. После каждой группы экспериментов на основе результатов всех предыдущих экспериментов (таблицы экспериментальных данных $v = \{x_{ij}\}$, $i = 1, \dots, t_\varphi$,

$j = 1, \dots, n$, $t_\varphi = \sum_{v=1}^{\varphi} r^{(v)}$) вводится функция распределения вероятностей $p_\varphi(x)$. Величина $p_\varphi(x)$ отражает вероятность того, что значение функции в точке x , если в ней провести эксперимент, будет наибольшим из всех значений функции, полученных за $(\varphi + 1)$ групп.

При введении функции $p_\varphi(x)$ используется следующая гипотеза: считается, что вероятность $p_\varphi(x)$ в окрестностях точек, в которых уже получено значение функции, тем больше, чем больше значение функции. Степень этого предположения зависит от числа проведенных экспериментов. Это делается следующим образом. Вычисляем энтропию

$$H^{(\varphi)} = - \sum_{x \in D} p_\varphi(x) \log p_\varphi(x), \quad \varphi = 0, 1, \dots, R.$$

В начале поиска, когда нет информации о предпочтении одних элементов множества D другим, вводится равномерное распределение вероятностей $p_0(x) = \frac{1}{N_0}$, где $N_0 = |D|$. Энтропия в этом случае максимальна и равна $H^{(0)} = \log N_0$. Задавая величину $H^{(\varphi)}$ ($\varphi = 0, 1, \dots, R$) как некоторую монотонно убывающую функцию $H(t_\varphi, k)$ от числа про-

веденных экспериментов $t_\varphi = \sum_{v=1}^{\varphi} r^{(v)}$, получаем сужение области по-

иска по мере проведения экспериментов. Скорость убывания этой функции задается некоторым параметром k , названным коэффициентом адаптации. Чем больше значение параметра k (больше степень адап-

таким), тем сильнее ограничение на класс рассматриваемых функций $\{f(x)\}$.

Расстановка $(\varphi+1)$ -й группы экспериментов делается в соответствии с функцией $p_\varphi(x)$. Будем считать, что расстановка $r^{(\varphi+1)}$ экспериментов на множестве D произведена в соответствии с функцией $p_\varphi(x)$, если для произвольного подмножества $D' \subseteq D$ число экспериментов пропорционально величине $\sum_{x \in D'} p_\varphi(x)$. По-видимому, выполнение

этого условия дает наилучшую расстановку экспериментов.

На основе сформулированного выше подхода были предложены адаптивные алгоритмы поиска наилучших значений функций для случая булевых переменных (алгоритм СПА [6]), переменных, замеряемых в шкале наименований [27,40], количественных переменных [26], разнотипных переменных [28]. В последней работе также приводится обзор указанных алгоритмов. Алгоритм СПА использовался для выбора наиболее информативной подсистемы признаков при решении задач из области медицины [29-33], из области социологии [34], экономики [35] и т.д.

Алгоритм поиска приближенного значения экстремума функции для разнотипных переменных (алгоритм ОПР) состоит из трех блоков: блока планирования группы экспериментов, блока обнаружения логических закономерностей на эмпирических таблицах и блока адаптации. Рассмотрим кратко указанные блоки.

Алгоритм планирования группы из r экспериментов предварительно будет описан для случая, когда число значений каждой переменной равно 1, а функция распределения вероятностей $p(x) = \frac{1}{1^n}$.

В дальнейшем точку n -мерного пространства будем называть реализацией. Из множества возможных реализаций, число которых равно 1^n , необходимо выбрать такое подмножество A из r реализаций ($r < 1^n$), чтобы для любой выбранной подсистемы из m переменных осуществлялся полный перебор всевозможных сочетаний значений из этих переменных. По-видимому, такой выбор подмножества из r реализаций достаточно полно удовлетворяет введенному равномерному распределению $p(x)$. Однако с увеличением числа m растет сложность алгоритма такого выбора и растет число реализаций r . Предложен выбор подмножества A для $m = 2$.

Пусть с помощью алгоритма планирования экспериментов получена таблица $\{x_{ij}\}$ ($i = 1, \dots, r$; $j = 1, \dots, n$; r - число реализа -

ций в группе; n - число переменных) и набор значений функции $\{x_{10}\}$, где $x_{10} = f(x_1)$. Далее для поиска логических закономерностей используется алгоритм ПИНЧ (см. §6).

Изменение функции распределения вероятностей $p_\phi(x) = \prod_{j=1}^n p_\phi(x_j)$

по мере проведения поиска (блок адаптации) подробно описан в работе [40]. Для того чтобы приписать вероятности $p_\phi(x_j)$ различным значениям переменной X_j , необходимо установить порядок их предпочтения в данный момент поиска. В рассматриваемом случае этот порядок определяется на основе набора логических закономерностей $\{a_v^u\}$ ($u = 1, \dots, l_j$; $v = 1, \dots, M_u$; l_j - число значений признака X_j ; M_u - число закономерностей, характеризующих выбранный на ϕ -м шаге адаптации интервал ϕ). Правило установления порядка на значениях переменной X_j приводится в работе [28].

§10. Применение алгоритмов обработки эмпирических таблиц

Рассмотренные в данной работе алгоритмы были использованы для решения задач из области медицины, социологии, экономики, геологии, химии, геофизики, а также ряда задач из области угольной промышленности. Ниже приводится краткая информация о решенных задачах.

1. В 1971 - 1974 гг. сотрудниками Новосибирского медицинского института, Института горного дела СО АН СССР и Института математики СО АН СССР проводилось крупное исследование по изучению вибрационной патологии среди рабочих завода им. В.П.Чкалова для разработки рекомендаций по ее профилактике, ранней диагностике и лечению больных. Проблема профилактики вибрационной болезни имеет большое социально-экономическое значение. Работа проводилась под руководством академика АМН СССР В.П.Кавнячева. Каждому рабочему ставилось в соответствие около 600 разнотипных признаков, характеризующих его физиологическое состояние и условия его труда. Число обследованных рабочих составляло около 900. Для статистического анализа использовался алгоритм КОРАЛЛ. Результаты обработки приводятся в коллективной монографии [31] и в работах [32, 33].

2. Алгоритм КОРАЛЛ использовался для построения логического решающего правила распознавания "протонных" и "непротонных" вспышек на Солнце. Работа проводилась совместно с СИБКИМРОм Академии наук. Каждая предвспышечная ситуация характеризовалась 24 разнотипными признаками. Наряду с количественными признаками (например,

скорости изменения площади группы пятен на Солнце) были признаки, замеренные в шкале наименований (тип группы пятен). Прогнозы "протонных" вспышек важны, например, для радиосвязи, при планировании времени запуска космических кораблей. Исходный материал состоял из описания 123 активных областей Солнца. Были получены решающие правила, позволявшие правильно классифицировать 48 из 56 предвсказанных состояний.

3. Алгоритм ПИНЧ, предназначенный для прогнозирования значения количественной переменной на основе логического решающего правила, использовался для прогноза горнотехнических параметров крепи проектируемых выработок в шахтах производственного объединения "Кузбассуголь". Прогнозирование основывалось на признаках, характеризующих каротажные кривые (кривые сопротивления горных пород, кривые искусственной и естественной радиоактивности). На основе полученных результатов Кузбасским политехническим институтом было разработано "Методическое пособие по проектированию типов и параметров крепи выемочных штреков шахт Кузнецкого бассейна". Внедрение указанного пособия дало значительный экономический эффект.

4. Кузбасским политехническим институтом проводилась работа по анализу пожароопасных шахт Кузбасса. Анализировался материал по 24 шахтам за 7 лет. На основе сформулированного в §5 подхода по упорядочению объектов по их перспективности была выработана методика прогноза и управления степенью эндогенной пожароопасности шахт. Практическая проверка методики на выемочных полях одной из шахт Кузбасса показала снижение уровня суммарных затрат на профилактику и тушение пожаров в среднем на 20%, что дает значительный экономический эффект.

5. Тот же подход был использован для упорядочения четырехсот геологических объектов (участков Приморского края) по их перспективности с точки зрения вероятности обнаружения на этих участках месторождения [21]. На контрольной выборке было показано, что упорядочение позволяет не проводить буровых работ примерно на 50% бесперспективных участках при условии, что при этом не будет пропущено ни одно месторождение. Кроме того, упорядочение объектов дает возможность оптимизировать поисково-разведочные работы.

6. Метод автоматической группировки объектов (§7) использовался, например, для классификации рабочих по их профориентации на основе разнотипных признаков, включающих в себя как медико-биологические признаки, так и признаки, характеризующие профессии.

7. Программа оптимизации ОПР была использована для решения задачи выбора наилучших значений десяти разнотипных переменных, задающих технологический процесс осаждения специального сплава при восстановлении крупногабаритных деталей сельскохозяйственной техники [36]. Применение классических методов оптимизации оказалось затруднительным, поскольку наряду с количественными переменными были и переменные, замеренные в шкале наименований. Были сформулированы условия, которые обеспечивают высокие физико-механические свойства осадка.

Алгоритм ОПР использовался также и для решения других задач [37].

§II. Пакет прикладных программ ОТЖС

Приведенные примеры использования алгоритмов обработки эмпирических таблиц показывают широкую их применимость в различных областях, а также достаточно большую экономическую эффективность от их внедрения. Для уменьшения времени между окончанием научных разработок и внедрением их в практику наиболее удобной формой их представления являются пакеты прикладных программ. Действительно, для широкого пользователя (медика, геолога, социолога и т.д.), для которого выбор наиболее подходящей программы в зависимости от эмпирической ситуации является затруднительным (для этого требуются достаточно полные знания о каждом методе, включенном в пакет), необходимо обеспечить автоматический режим вычислений по достаточно простому запросу.

Естественно предложить для описания эмпирической ситуации, возникающей при обработке таблиц экспериментальных данных, использовать следующие семь параметров: 1) тип задачи (распознавание образов, автоматическая группировка объектов, упорядочивание объектов, оптимизация и т.д.); 2) тип шкалы (все признаки исходной системы замерены только в шкале наименований, только в шкале порядка, только в шкале интервалов и отношений либо в разных шкалах); 3) число признаков; 4) число реализаций; 5) наличие пропусков в таблицах; 6) характеристика зависимости признаков; 7) ограничение на форму таксона. Набор значений этих параметров образует то или иное управляющее слово (запрос), на основе которого соответствие с предложенной таблицей решений автоматически выбирается соответствующая данной эмпирической ситуации программа. Таблица ре-

шений изменяется по мере решения задач на основе экспертных оценок эффективности программ, включенных в пакет.

В Институте математики СО АН СССР и Новосибирском государственном университете создан пакет прикладных программ ОТЭС для ЭВМ "Минск-32" и ЕС ЭВМ [38,39], в который вошли в качестве составной части программы, реализующие рассмотренные в данной работе алгоритмы. Пакет ОТЭС был реализован в рамках работ по созданию пакетов прикладных программ для распознавания и автоматической классификации по плану Госкомитета по науке и технике и является одним из результатов работы рабочей группы РГ-2 Комиссии многостороннего сотрудничества Академий наук социалистических стран по проблеме "Научные вопросы вычислительной техники".

Л и т е р а т у р а

1. ПАНЦАГЛЬ И. Теория измерений. М., "Мир", 1976.
2. ЛЕОВ Г.С., КОТЮКОВ В.И., МАНОХИН А.Н. Об одном алгоритме распознавания в пространстве разнотипных признаков. -В кн.: Вычислительные системы. Вып. 55, Новосибирск, 1973, с. 98-107.
3. ВАПНИК В.Н., ЧЕРВОНЕНКИС А.Я. Теория распознавания образов. М., "Наука", 1974.
4. ЛЕОВ Г.С., МАНОХИН А.Н. Распознавание образов при разнотипных признаках в условиях малой выборки. -В кн.: Статистические проблемы управления (материалы Всесоюзного семинара "Проблемы малых выборок в распознавании образов"). Вып. 14, Вильнюс, 1976, с. 57-63.
5. ЛЕОВ Г.С., МАНОХИН А.Н., МАШАРОВ Ю.П. Логические решающие правила в распознавании образов. -В кн.: Машинные методы обнаружения закономерностей. Новосибирск, 1976, с.64-74.
6. ЛЕОВ Г.С. Выбор эффективной системы зависимых признаков. -В кн.: Вычислительные системы. Вып. 19, Новосибирск, 1965, с.21-34.
7. АЙЗЕРМАН М.А., БРАВЕРМАН Э.М., РОЗОНДЭР Л.И. Метод потенциальных функций в теории обучения машин. М., "Наука", 1970.
8. ДУДА Р., ХАРТ П. Распознавание образов и анализ сцен. М., "Мир", 1976.
9. НИЛЬСЕН Н. Обучающиеся машины. М., "Мир", 1967.
10. ЕЛКИНА В.Н., ЗАГОРУЙКО Н.Г. Количественные критерии качества таксономии и их использование в процессе принятия решений. -В кн.: Вычислительные системы. Вып. 36. Новосибирск, 1969, с.29-46.
11. ЗАГОРУЙКО Н.Г. Методы распознавания и их применение. М., "Сов.радио", 1972.
12. ВОРОНИН Ю.А. Введение мер сходства и связи для решения геолого-географических задач. -Докл. АН СССР, т.199, № 5, 1971, с. 1011-1014.

13. ЖУРАВЛЕВ Ю.И., НИКИФОРОВ В.В. Алгоритмы распознавания, основанные на вычислении оценок. - "Кибернетика", № 3, 1971, с.1-12.
14. ВАЙНЦВАЙГ М.Н. Алгоритмы обучения распознаванию образов. "Кора". - В кн.: Алгоритмы обучения распознаванию образов. М., "Сов. радио", 1973, с.110-115.
15. ЖУРАВЛЕВ Ю.И., ДМИТРИЕВ А.Н., КРЕНДЕЛЕВ Ф.П. О математических принципах классификации предметов и явлений. - В кн.: Дискретный анализ. Новосибирск, №7, 1966, с.3-15.
16. ГОРЕЛИК А.Л., СКРИПНИК В.А. Методы распознавания. М., "Высшая школа", 1977.
17. РАУДИС Ш. Ограниченность выборки в задачах классификации. - В кн.: Статистические проблемы управления, вып. 18, Вильнюс, 1976.
18. ЛБОВ Г.С. О представительности выборки при выборе эффективной системы признаков. - В кн.: Вычислительные системы. Вып.22, Новосибирск, 1966, с. 39-58.
19. ЛБОВ Г.С., КОТЮКОВ В.И., МАШАРОВ Ю.П. Метод обнаружения логических закономерностей на эмпирических таблицах. - В кн.: Эмпирическое предсказание и распознавание образов. (Вычислительные системы, вып.67.) Новосибирск, 1976, с. 29-42.
20. МАНОХИН А.Н. Методы распознавания образов, основанные на логических решающих функциях. - В кн.: Эмпирическое предсказание и распознавание образов. (Вычислительные системы, вып. 67.) Новосибирск, 1976, с. 42-53.
21. ЛБОВ Г.С., МАНОХИН А.Н. О решении задачи прогнозирования месторождений. Отчет Института математики СО АН СССР, Новосибирск, 1976.
22. МАНОХИН А.Н. Об одном подходе к прогнозированию перспективности объектов. - В кн.: Методы обработки информации. (Вычислительные системы, вып. 74.) Новосибирск, 1976, с. 108-128.
23. ЛБОВ Г.С. Об одном непараметрическом подходе в задачах эмпирического предсказания. - В кн.: Адаптивные системы и их приложения, Новосибирск, "Наука", 1976, с.78-81.
24. ЛБОВ Г.С., МАНОХИН А.Н. Логические и линейно-логические функции оценивания. - В кн.: I Областная научно-практическая конференция по надежности научно-технических прогнозов. Новосибирск, 1978, с. 58-59.
25. ЛБОВ Г.С., МАШАРОВ Ю.П. Метод динамического прогнозирования, использующий логические решающие правила. - Настоящий сборник, с. 65-74.
26. ЛБОВ Г.С., ТРУНОВ А.А. Об одном алгоритме поиска глобального экстремума функции. - В кн.: Эмпирическое предсказание и распознавание образов. (Вычислительные системы, вып. 67.) Новосибирск, 1976, с. 29-41.
27. LBOV G.S. Training for extremum determination of function of variables measured in names scale. - In: Sec.Int.Conf. on Artificial intelligence, London, 1972, p.418-423.
28. ЛБОВ Г.С. Алгоритмы поиска приближенного значения глобального экстремума функции. - В кн.: Проблемы случайного поиска. Вып.8, Рига, "Зинатне", 1978.

29. БОРОВКОВ А.А., ЗАСЛАВСКИЙ А.Е., ЛБОВ Г.С., СИЧЕВА Н.М. Анализ багнестокардиограмм в целях диагностики заболеваний сердца. -В кн.: Физико-математические методы исследований в биологии и медицине. (Материалы к первой Новосибирской конференции). Новосибирск, 1965, с. 58-62.

30. АБДУЛЛАЕВА Н.С., ЛБОВ Г.С. Выбор существенных признаков для диагностики врожденных пороков сердца. В кн.: Вопросы кибернетики, вып. 51, Ташкент, "ФАН", с.88-93.

31. Проблемы охраны труда при клепальных работах. Под ред. Белевонской Н.П., М., "Наука", 1978.

32. ЛБОВ Г.С., КОТКОВ В.И. Методы распознавания образов в комплексных медицинских исследованиях. -Биологическая и медицинская кибернетика, ч.4, Медицинская кибернетика, М., 1974.

33. LBOV G.S. Application of pattern recognition method vi-viopathology problems.-EunHore international ISME Symposium, Tokyo, 1977.

34. Распознавание образов в социальных исследованиях. Под ред. Загоруйко Н.Г. и Заславской Т.И., Новосибирск, "Наука", 1968.

35. РОЗИН Б.Б., ЛБОВ Г.С., КУРАВЕЛЬ Н.М. Использование методов распознавания образов в технико-экономическом анализе производства. -В кн.: Проблемы экономико-статистического анализа. Новосибирск, "Наука", 1969, с. 239-252.

36. РЕВЯКИН В.П., ЛБОВ Г.С., ШИПКИН Г.М. Исследование процесса электролитического натирания твердыми сплавами методом случайного поиска с адаптацией. -В кн.: Электрохимия и ее применение в промышленности. Тюмень, 1971, с. 31-38.

37. ЛБОВ Г.С. Об оптимизации функционирования измерительных гидроакустических систем. -Труды I-й Всесоюзной школы-семинара по статистической гидроакустике, Новосибирск, "Наука", 1970, с. 271-282.

38. ЗАГОРУЙКО Н.Г., ЛБОВ Г.С., МАШАРОВ Ю.П. Пакет прикладных программ для обработки таблиц экспериментальных данных ОТЗКС. -В сб.: Вопросы кибернетики. Материалы III Ленинградского симпозиума "Теория адаптивных систем", ч.2, М., 1977, с. 5-9.

39. Пакет прикладных программ для обработки таблиц экспериментальных данных ОТЗКС. Новосибирск, Изд-во НГУ, 1977.

40. ЛБОВ Г.С. Адаптивный поиск экстремума функций от переменных, измеренных в шкале наименований. -В сб.: Вычислительные системы. Вып. 44. Новосибирск, 1971, с. 13-22.

41. ТРУНОВ А.А. Алгоритм построения кусочно-линейных логических правил правил распознавания. Материалы XII областной научно-технической конференции, посвященной дню радио и связи. Новосибирск, 1978, с.43.

Поступила в ред.-изд.отд.
4 августа 1978 года