

Математические методы исследования

УДК 519.24

ТРИ ОСНОВНЫХ РЕЗУЛЬТАТА МАТЕМАТИЧЕСКОЙ ТЕОРИИ КЛАССИФИКАЦИИ

© А. И. Орлов¹*Статья поступила 22 июня 2015 г.*

Математическая теория классификации весьма многообразна, содержит большое число подходов, моделей, методов, алгоритмов. Выделим в ней три результата — оптимальный метод диагностики (дискриминантного анализа), адекватный показатель качества алгоритма дискриминантного анализа, утверждение об остановке после конечного числа шагов итерационных алгоритмов кластер-анализа. На основе леммы Неймана — Пирсона показано, что оптимальный метод диагностики существует и выражается через плотности распределения вероятностей, соответствующие классам. Если плотности неизвестны, следует использовать их непараметрические оценки по обучающим выборкам. Часто используют такой показатель качества алгоритма диагностики, как «вероятность (или доля) правильной классификации (диагностики)»; чем этот показатель больше, тем алгоритм лучше. Показана нецелесообразность повсеместного применения этого показателя и обоснован другой — «прогностическая сила», полученная путем пересчета на модель линейного дискриминантного анализа. Остановка после конечного числа шагов итерационных алгоритмов кластер-анализа продемонстрирована на примере метода k -средних. Эти результаты являются важными в теории классификации, они должны быть известны каждому специалисту, развивающему эту теорию или применяющему ее.

Ключевые слова: математическая теория классификации; математическая статистика; прикладная статистика; диагностика; дискриминантный анализ; лемма Неймана — Пирсона; показатель качества алгоритма диагностики; вероятность правильной классификации; прогностическая сила; кластер-анализ; остановка итерационного алгоритма; метод k -средних.

Методы классификации — неотъемлемая часть математических методов исследования, интересная теоретически и важна практически. Обзоры этой научной области представлены в работах [1–3]. Многие математические методы классификации относятся к непараметрической статистике [4] и к нечисловой статистике [5], т.е. являются неотъемлемой составной частью основного потока современных научных исследований, порожденных новой парадигмой прикладной статистики [6].

В многообразии результатов математической теории классификации выделим три: оптимальный метод диагностики (дискриминантного анализа), адекватный показатель качества алгоритма дискриминантного анализа, доказательство сходимости итерационных алгоритмов кластер-анализа. По нашей оценке, эти результаты являются основными в теории классификации.

Оптимальный метод диагностики

Рассмотрим задачу диагностики с двумя классами. Решение принимают по основе значения x — эле-

мента некоторого пространства. Элементы первого класса имеют плотность $f(x)$, элементы второго — плотность $g(x)$. Поступает на рассмотрение новый объект со значением X . К какому классу его отнести?

Задачу диагностики можно переформулировать в терминах теории проверки статистических гипотез. Пусть согласно нулевой гипотезе H_0 результат наблюдения X имеет распределение с плотностью $f(x)$, а согласно альтернативной гипотезе H_1 — с плотностью $g(x)$. Отнесение X к первому классу соответствует принятию гипотезы H_0 (и отклонению гипотезы H_1), а отнесение X ко второму классу — принятию гипотезы H_1 (и отклонению гипотезы H_0).

В теории проверки статистических гипотез выявлена важная роль критерия отношения правдоподобия (см., например, [7]). Статистика этого критерия имеет вид

$$Q(x) = f(x)/g(x). \quad (1)$$

Правило принятия решения основано на сравнении с порогом C значения статистики критерия $Q(X)$,

¹ Институт высоких статистических технологий и эконометрики Московского государственного технического университета им. Н. Э. Баумана, Москва, Россия; Московский физико-технический институт, г. Долгопрудный Московской обл., Россия; Центральный научно-исследовательский институт машиностроения, г. Королёв Московской области, Россия; e-mail: prof-orlov@mail.ru

рассчитанного для поступившего на рассмотрение нового объекта со значением X . Таким образом, если $Q(X) > C$, то X относят к первому классу, в противном случае — ко второму.

С точки зрения здравого смысла критерий отношения правдоподобия является естественным, как отношение шансов (вероятностей) на то (того), что новый объект со значением X относится к первому или ко второму классу соответственно. Важно, что согласно лемме Неймана – Пирсона этот критерий является наиболее мощным среди всех статистических критериев, имеющих один и тот же заданный уровень значимости (понятия «уровень значимости» и «мощность критерия» — базовые в математической статистике). (Строго говоря, под термином «лемма» понимают верное или доказанное утверждение, полезное не само по себе, а для доказательства других утверждений. Однако лемма Неймана – Пирсона — основной результат математической статистики, важный сам по себе. Поэтому лемму Неймана – Пирсона часто называют фундаментальной леммой математической статистики.)

Итак, оптимальный метод диагностики существует и задается с помощью статистики $Q(X)$ (см. формулу (1)).

Однако при решении практических задач диагностики плотности $f(x)$ и $g(x)$ обычно неизвестны. В таких случаях строят правило диагностики на основе обучающих выборок. А именно, предполагают, что имеются m объектов из первого класса (обучающая выборка для первого класса) и n объектов из второго класса (обучающая выборка для второго класса). В вероятностно-статистической теории принимают, что обучающую выборку можно моделировать как совокупность независимых одинаково распределенных случайных объектов с соответствующей плотностью. Развита непараметрические методы состоятельного оценивания неизвестной плотности [8, 9]. Пусть $f_m(x)$ и $g_n(x)$ — состоятельные оценки плотностей $f(x)$ и $g(x)$ соответственно по обучающим выборкам. Рассмотрим выборочный аналог статистики критерия отношения правдоподобия

$$Q_{mn}(x) = f_m(x)/g_n(x). \quad (2)$$

Из состоятельности $f_m(x)$ и $g_n(x)$ следует, что $Q_{mn}(x)$ для того же элемента x является состоятельной оценкой $Q(x)$ при безграничном росте объемов обучающих выборок. При справедливости обычно выполненного предположения равномерной сходимости из оптимальности критерия отношения правдоподобия для полностью известных плотностей вытекает асимптотическая оптимальность выборочного аналога этого критерия, основанного на сравнении с порогом C значения статистики (2).

В задачах диагностики со многими классами оптимальное решение также выражается через плотности, соответствующие классам, например, при постановке задачи в терминах статистических реша-

ющих правил [10, 11]. Во всех таких случаях асимптотически оптимальное решение получаем путем замены неизвестных плотностей их состоятельными оценками [8, 9].

Наличие описанных выше оптимальных и асимптотически оптимальных правил диагностики (дискриминантного анализа, распознавания образов с учителем) не означает, что не следует разрабатывать новые алгоритмы диагностики, например, исходя из необходимости сокращения машинной памяти и времени на расчеты. На наш взгляд, *следует сравнивать новые алгоритмы с известными оптимальными и асимптотически оптимальными алгоритмами* по тем или иным показателям качества.

Прогностическая сила — адекватный показатель качества алгоритма диагностики

Часто используют такой показатель качества алгоритма диагностики, как «вероятность (или доля) правильной классификации (диагностики)» [12, 13]; чем этот показатель больше, тем алгоритм лучше. Покажем нецелесообразность повсеместного применения этого показателя и обоснуем другой — «прогностическую силу», найденную путем пересчета на модель линейного дискриминантного анализа.

Приведем используемую в дальнейшем вероятностно-статистическую модель диагностики. Пусть классифицируемые объекты описываются переменными x , лежащими в некотором пространстве Z , два класса — это два распределения вероятностей с плотностями $f(x)$ и $g(x)$, $x \in Z$, соответственно; если Z дискретно, то под $f(x)$ и $g(x)$ понимаем не плотности, а вероятности попадания в точку x .

Типичная схема разработки конкретного математического метода диагностики такова. С помощью специалистов соответствующей прикладной области составляют две обучающие выборки — объема m_0 из первого класса и объема n_0 из второго класса. На их основе определяют решающее правило $A: Z \rightarrow \{1, 2\}$, ставящее в соответствие результату наблюдения $x \in Z$ номер класса, к которому его следует отнести. Качество работы алгоритма проверяют по контрольной выборке, состоящей из m элементов первого класса и n элементов второго. Результаты проверки удобно записать в виде, представленном в таблице. Ясно, что о качестве алгоритма диагностики, т.е. решающего правила A , надо судить на основе этих результатов. Естественно использовать $\kappa = a/m$ и $\lambda = d/n$ — доли правильной диагностики в первом и во втором классах.

Доля правильной диагностики

$$\mu = \frac{a+d}{m+n} = \pi_1 \kappa + \pi_2 \lambda,$$

где $\pi_1 = m/(m+n)$ и $\pi_2 = n/(m+n)$ — априорные доли первого и второго классов. Очевидно, $\pi_1 + \pi_2 = 1$.

В вероятностной модели речь идет об априорных вероятностях классов. Удобно и в статистической (см. таблицу), и в вероятностной постановках использовать единый термин «доля», поскольку это не приводит к недоразумениям.

Рассмотрим сначала случай, когда $\min(m_0, n_0) \rightarrow \infty$, т.е. плотности f и g можно считать известными. При применении теории статистических решений к задачам диагностики [13] считаются известными также априорные вероятности классов и потери $C(j|i)$ от ошибочной диагностики — от отнесения объекта i -го класса к классу с номером j . Доказано (при любом числе классов), что в случае, когда все потери $C(j|i)$ равны между собой, минимизация суммарных потерь эквивалентна максимизации вероятности правильной диагностики [13], состоятельной оценкой которой служит μ . Видимо, этот факт вместе с вычислительной простотой и наглядностью показателя μ является причиной широкого использования доли (вероятности) правильной классификации μ как показателя качества алгоритма диагностики.

Оптимальное решающее правило при совпадающих потерях $C(1|2) = C(2|1)$ таково: если

$$\frac{f(x)}{g(x)} > \frac{\pi_2(\infty)}{\pi_1(\infty)}, \quad (3)$$

то x относят к первому классу, в противном случае — ко второму; здесь $\pi_1(\infty)$ и $\pi_2(\infty)$ — априорные вероятности первого и второго классов соответственно. Как показано выше, если исходить из другой оптимизационной постановки (из минимизации ошибки второго рода при ограничении на уровень значимости), лемма Неймана – Пирсона также дает правило, основанное на отношении плотностей вероятностей. В большинстве прикладных задач плотности вероятности неизвестны. Однако при $\min(m_0, n_0) \rightarrow \infty$ их можно заменить состоятельными оценками, например, непараметрическими ядерными оценками плотности [8, 9, 14], и получить *универсальное асимптотически оптимальное правило диагностики*, с которым необходимо сравнивать все другие правила диагностики [1, 15]. Выбор других правил диагностики для решения конкретных прикладных задач должен быть обоснован соответствующими задачами критериями, например, быстродействия алгоритмов и объемов имеющейся информации [16].

Обратим внимание, что в случае, когда априорная вероятность одного из классов существенно больше априорной вероятности другого, правило (3) может любое наблюдение относить к классу с наибольшей априорной вероятностью.

Разберем ситуацию подробнее. Пусть имеется некоторый алгоритм диагностики на два класса с долями правильной диагностики κ — в первом классе и λ — во втором. Сравним его с двумя тривиальными алгоритмами диагностики. Первый тривиальный алгоритм относит все классифицируемые объекты к первому

классу, для него $\kappa = 1$ и $\lambda = 0$, следовательно, $\mu = \pi_1$. Второй тривиальный алгоритм относит все классифицируемые объекты ко второму классу, для него $\kappa = 0$ и $\lambda = 1$, следовательно, $\mu = \pi_2$.

В качестве показателя качества алгоритма диагностики используем долю правильной диагностики μ . Когда первый тривиальный алгоритм лучше исходного? Когда $\pi_1 > \kappa\pi_1 + \lambda\pi_2$, т.е.

$$\frac{\lambda}{1 - \kappa + \lambda} < \pi_1$$

(с учетом того, что $\pi_1 + \pi_2 = 1$). Когда второй тривиальный алгоритм лучше исходного? Когда $\pi_2 > \kappa\pi_1 + \lambda\pi_2$, т.е.

$$\pi_1 < \frac{1 - \lambda}{1 + \kappa - \lambda}.$$

Таким образом, для любого заданного алгоритма диагностики существуют границы $d(1)$ и $d(2)$ для доли первого класса π_1 в объединенной контрольной выборке такие, что при $\pi_1 < d(1)$ рассматриваемый алгоритм хуже второго тривиального алгоритма, а при $\pi_1 > d(2)$ он хуже первого тривиального алгоритма.

Разобранная ситуация встречается на практике. В конкретной прикладной медицинской задаче величина μ оказалась больше для тривиального прогноза, согласно которому у всех больных течение заболевания (инфаркта миокарда) будет благоприятно. Тривиальный прогноз сравнивался с алгоритмом выделения больных с прогнозируемым тяжелым течением заболевания. Он был разработан группой математиков и кардиологов под руководством И. М. Гельфанда. Применение этого алгоритма с медицинской точки зрения вполне оправдано [17–19]. Итак, по доле правильной классификации μ алгоритм группы И. М. Гельфанда оказался хуже тривиального — объявить всех пациентов легкобольными, т.е. не требующими специального наблюдения. Этот вывод очевидно нелеп. И причина появления нелепости вполне понятна. Хотя доля тяжелобольных невелика, но смертельные исходы сосредоточены именно в этой группе больных. Поэтому целесообразна *гипердиагностика* — рациональнее часть легкобольных объявить тяжелобольными, чем сделать ошибку в противоположную сторону.

Поэтому мы полагаем, что использовать в качестве показателя качества алгоритма диагностики долю правильной диагностики μ нецелесообразно.

Работы группы И. М. Гельфанда [17–19] показывают также, что теория статистических решений не может быть основой для выбора показателя качества диагностики, применение ее требует знания потерь от ошибочной диагностики. В большинстве же научно-технических и экономических задач определить потери, как уже отмечалось, сложно, в частности, из-за необходимости оценивать человеческую жизнь в денежных единицах. По этическим соображениям это, на наш взгляд, недопустимо. Сказанное не означает отри-

цания пользы страхования, но, очевидно, страховые выплаты следует рассматривать лишь как способ первоначального смягчения потерь от утраты близких. Следовательно, применение теории статистических решений в рассматриваемой постановке вряд ли возможно, поскольку оценить количественно потери от смерти больного нельзя по этическим соображениям.

С целью поиска приемлемого показателя качества диагностики рассмотрим [20] широко известную параметрическую вероятностную модель (модель линейного дискриминантного анализа), в которой Z — конечномерное пространство, $f(x)$ и $g(x)$ — многомерные нормальные плотности с математическими ожиданиями m_1 и m_2 соответственно и совпадающими ковариационными матрицами Σ . Тогда при произвольных априорных вероятностях и потерях оптимальное решающее правило определяется плоскостью

$$H(x) = \left(x - \frac{m_1 + m_2}{2} \right)^T \Sigma^{-1} (m_1 - m_2) = C, \quad x \in Z,$$

где константа C зависит от априорных вероятностей классов $\pi_1(\infty)$ и $\pi_2(\infty)$, а также от потерь $C(1|2)$ и $C(2|1)$ [10, с. 186]. Основным параметром модели — расстояние Махаланобиса между классами

$$d = \{(m_1 - m_2)^T \Sigma^{-1} (m_1 - m_2)\}^{1/2}.$$

(Величину d^2 нельзя называть «расстоянием», как это делается в [10, с. 187], поскольку для d^2 не выполнено неравенство треугольника. Всем аксиомам, задающим метрику, удовлетворяет d .)

Через d и C/d выражаются вероятности правильной диагностики $\kappa(\infty)$ и $\lambda(\infty)$, являющиеся пределами при $\min(m, n) \rightarrow \infty$ ранее введенных долей κ и λ . Если X — случайная величина с описанной выше плотностью $f(x)$, то случайная величина $H(X)$ имеет нормальное распределение с математическим ожиданием $d^2/2$ и дисперсией d^2 . Аналогично для случайной величины Y с плотностью $g(x)$ случайная величина $H(Y)$ имеет нормальное распределение с математическим ожиданием $(-d^2/2)$ и дисперсией d^2 [10, с. 187]. Поэтому

$$\begin{aligned} \kappa(\infty) &= P[H(X) > C] = \Phi\left(\frac{d}{2} - \frac{C}{d}\right), \\ \lambda(\infty) &= P[H(Y) \leq C] = \Phi\left(\frac{d}{2} - \frac{C}{d}\right), \end{aligned} \quad (4)$$

где $\Phi(x)$ — функция стандартного нормального распределения с математическим ожиданием 0 и дисперсией 1. Из (4) следует, что при любых потерях $C(1|2)$ и

$C(2|1)$ и любых априорных вероятностях π_1 и π_2 при использовании оптимального решающего правила величина

$$d_0 = \Phi^{-1}(\kappa(\infty)) + \Phi^{-1}(\lambda(\infty))$$

постоянна и равна d , где $\Phi^{-1}(y)$ — функция, обратная к $\Phi(x)$. Следовательно, именно расстояние Махаланобиса d целесообразно рассматривать как меру различия между классами, заданными плотностями рассматриваемого вида. Чтобы выразить меру различия в тех же единицах, что и вероятности правильной диагностики, введем «прогностическую силу»

$$\delta = \Phi(d/2),$$

при $C = 0$ равную совпадающим значениям правильной диагностики $\kappa(\infty)$ и $\lambda(\infty)$ из (4).

Предлагаем в качестве показателя качества производного алгоритма диагностики использовать эмпирическую прогностическую силу

$$\delta^* = \Phi(d^*/2), \quad (5)$$

$$\text{где} \quad d^* = \Phi^{-1}(\kappa) + \Phi^{-1}(\lambda) \quad (6)$$

выборочная оценка расстояния Махаланобиса; доли κ и λ правильной диагностики в классах определены по данным таблицы. Таким образом, вместо взвешенного по априорным долям классов среднего арифметического μ долей κ и λ правильной диагностики в классах предлагаем использовать их среднее по Колмогорову с весовой функцией Φ^{-1} (о средних по Колмогорову см., например, [5]).

Если классы описываются выборками из многомерных нормальных совокупностей с одинаковыми матрицами ковариаций, а для классификации применяется классический линейный дискриминантный анализ Р. Фишера [10, 20], то величина d^* представляет собой состоятельную статистическую оценку расстояния Махаланобиса между двумя рассматриваемыми совокупностями, причем независимо от порогового значения, определяющего конкретное решающее правило. В общем случае показатель d^* вводится как эвристический.

Пример. Если $\kappa = 0,90$ и $\lambda = 0,80$, то $\Phi^{-1}(\kappa) = 1,28$ и $\Phi^{-1}(\lambda) = 0,84$, откуда $d^* = 2,12$ и эмпирическая прогностическая сила $\delta^* = \Phi^{-1}(1,06) = 0,86$. При этом доля правильной диагностики μ может принимать любые значения между 0,80 и 0,90, в зависимости от долей классов в объединенной совокупности π_i , $i = 1, 2$; $\pi_1 + \pi_2 = 1$.

Изучено асимптотическое распределение δ^* , разработаны методы расчета доверительных границ для прогностической силы по данным таблицы [21, 22].

Как проверить обоснованность применения прогностической силы, т.е. допустимость пересчета на модель линейного дискриминантного анализа? В ряде прикладных задач диагностики вычисляют значение некоторого прогностического индекса (фактора, пере-

Результаты работы алгоритма диагностики

Диагностируемые элементы	Число элементов	Отнесено к первому классу	Отнесено ко второму классу
Первого класса	m	a	b
Второго класса	n	c	d

менной) y и решение принимают на основе его сравнения с некоторым заданным порогом c . Объект относят к первому классу, если $y \leq c$, ко второму, если $y > c$. Прогностический индекс — это обычно линейная функция от характеристик рассматриваемых объектов, другими словами, от координат векторов, описывающих объекты. Возьмем два значения порога — c_1 и c_2 . Если классы описываются выборками из многомерных нормальных совокупностей с одинаковыми матрицами ковариаций, а для построения прогностического индекса применяется классический линейный дискриминантный анализ Р. Фишера, другими словами, если пересчет на модель линейного дискриминантного анализа обоснован, то, как можно показать, «прогностические силы» для обоих правил совпадают: $\delta(c_1) = \delta(c_2)$. Выполнение этого равенства можно проверить как статистическую гипотезу. Если эта гипотеза принимается, то целесообразность использования прогностической силы подтверждается, и есть основания значение $\delta^*(c_1) \approx \delta^*(c_2)$ рассматривать как объективную оценку качества алгоритма диагностики. Если же рассматриваемая гипотеза отклоняется, т.е. значения $\delta^*(c_1)$ и $\delta^*(c_2)$ сильно различаются, то пересчет на модель линейного дискриминантного анализа и использование расстояния Махаланобиса для измерения различия классов и прогностической силы как показателя качества диагностики некорректны. Способ проверки, т.е. соответствующий критерий проверки статистической гипотезы $\delta(c_1) = \delta(c_2)$, включая алгоритмы расчетов, разработан в [21, 22].

Сходимость итерационных алгоритмов кластер-анализа

Сначала обсудим один из широко применяемых методов кластер-анализа — метод k -средних. Он предназначен для разбиения исходного множества элементов (объектов или признаков) $x_1, x_2, \dots, x_n$, лежащих в некотором пространстве Z , на k кластеров. Опишем его в предлагаемой нами общей формулировке.

Метод основан на использовании функции $f: Z^2 \rightarrow [0, +\infty)$ в пространстве Z , имеющей смысл показателя различия (меры близости, расстояния), т.е. чем элементы x и y дальше отстоят друг от друга (в смысле, принятом в предметной области), тем $f(x, y)$ больше, причем $f(x, x) = 0$ для любого x из Z .

Метод k -средних — итерационный. Каждая итерация состоит из двух шагов — распределения элементов $x_1, x_2, \dots, x_n$ по k кластерам (первый шаг) и расчете центров кластеров (второй шаг).

Исходная информация перед началом первого шага — центры k кластеров, т.е. точки $a_1, a_2, \dots, a_k$ пространства Z . Для каждого из элементов $x_1, x_2, \dots, x_n$ находим ближайший центр. Для каждого из центров формируем кластер, состоящий из тех элементов $x_1, x_2, \dots, x_n$, которые ближе к этому центру, чем к другим центрам. Для строгости уточним: если некий эле-

мент находится на равном расстоянии от нескольких центров, то относим его к кластеру, центр которого имеет наименьший номер.

Исходная информация перед началом второго шага итерации — разбиение на k кластеров (полученное на первом шаге). На каждом шаге рассчитываем его центр. В соответствии с методологией статистики объектов нечисловой природы [5] в качестве центра кластера используем эмпирическое среднее элементов, включенных в кластер. Таким образом, для кластера A в качестве центра используем решение оптимизационной задачи

$$f(A, y) = \sum_{x \in A} f(x, y) \rightarrow \min_{y \in Z}. \quad (7)$$

Приведем примеры [23]. Если Z — конечномерное евклидово пространство, $f(x, y)$ — квадрат евклидова расстояния между точками x и y , то решением задачи (7) является центр тяжести точек, включенных в кластер. Другими словами, в этом пространстве надо взять значения координат точек, включенных в кластер, и рассчитать их среднее арифметическое. Это и будет значением первой координаты искомого центра. Если же $f(x, y)$ — блочное расстояние между точками x и y , то решением задачи (7) является точка, каждая координата которой — выборочная медиана для точек, включенных в кластер. Если Z — то или иное пространство бинарных отношений, $f(x, y)$ — расстояние Кемени, то решением задачи (7) является медиана Кемени.

Замечание. Если $Z = R^k$ — конечномерное евклидово пространство, то его точки можно обозначить $x = (x_1, x_2, \dots, x_k)$ и $y = (y_1, y_2, \dots, y_k)$. Блочное расстояние между точками x и y имеет вид

$$f(x, y) = \sum_{1 \leq i \leq k} |x_i - y_i|$$

Если Z — то или иное пространство бинарных отношений, то его точки описываются некоторыми матрицами из 0 и 1 заданного порядка k . Их можно обозначить $x = |x_{ij}|$ и $y = |y_{ij}|$. Расстояние Кемени между точками x и y имеет вид

$$f(x, y) = \sum_{1 \leq i, j \leq k} |x_{ij} - y_{ij}|$$

Условия существования решения задачи (7) найдены в статистике объектов нечисловой природы [5, 23]. Если решение задачи (7) не единственно, то правила однозначного выбора центра кластера должны быть специально указаны. Здесь нет необходимости останавливаться на этих подробностях.

Итогом второго шага итерации являются (новые) центры k кластеров.

Переходим к следующей итерации. Строим (новые) кластеры. Распределение элементов по (новым) кластерам, вообще говоря, изменится (по сравнению с распределением в начале предыдущей итерации). Находим для (новых) кластеров центры согласно формуле (7). Затем — следующий шаг.

Для запуска алгоритма необходимо перед началом первой итерации тем или иным способом задать центры k кластеров. Можно взять первые k из элементов $x_1, x_2, \dots, x_n$, или случайно выбрать k из них, или задать k точек из пространства Z (например, стараясь равномерно охватить естественную область изменения подлежащих кластеризации элементов), и т.д. Влияние начального задания элементов уменьшается при увеличении числа итераций.

Проблема сходимости итерационного алгоритма k -средних такова: остановится ли процесс итераций? Остановка возможна, если для некоторой итерации «новые» кластеры совпадут со «старыми», тогда и «новые» центры совпадут со «старыми».

Априори есть две возможности: либо итерационный алгоритм остановится, дав точное решение задачи кластер-анализа в рассматриваемой обстановке, либо итерации могут продолжаться бесконечно и тогда для получения приближенного решения надо тем или иным способом его останавливать.

Ответ на поставленный вопрос может быть получен двумя способами — экспериментальным (путем проведения расчетов по рассматриваемому алгоритму при тех или иных исходных данных) или теоретическим (путем доказательства соответствующих теорем).

Во всех проведенных расчетах итерационный алгоритм метода k -средних останавливался, т.е. эксперименты обосновывают предположение о том, что итерационный алгоритм остановится. Но ограниченный прошлый опыт не дает гарантий на будущее, с его помощью можно получить только правдоподобное утверждение. Для полной гарантии необходимы строгие утверждения — теоремы. И они были получены [24 – 26].

Поскольку строгие формулировки и доказательства громоздки, поясним здесь лишь основную идею. Выявим причину, по которой рассматриваемый алгоритм всегда останавливается.

Разбиение исходного множества элементов $x_1, x_2, \dots, x_n$, на кластеры, полученное на итерации с номером t , обозначим $\Psi = \{A_1, A_2, \dots, A_k\} = \{A_1(t), A_2(t), \dots, A_k(t)\}$. Рассмотрим один из основных показателей качества кластеризации — внутрикластерный разброс

$$g(\Psi) = \sum_{i=1}^k f(A_i, y_i), \quad (8)$$

где функция $f(A, y)$ определяется формулой (7); $y_i = y_i(t)$ — центр кластера A_i , полученный на итерации с номером t путем решения оптимизационной задачи (7), $i = 1, 2, \dots, k$. На втором шаге очередной итерации каждое слагаемое в правой части (8) либо уменьшается (если новый центр отличен от предыдущего), либо остается тем же самым (если новый центр совпадает с предыдущим). Следовательно, и сумма этих слагаемых — внутрикластерный разброс $g(\Psi)$ —

либо уменьшается, либо остается тем же самым. При этом $g(\Psi)$ не меняется тогда и только тогда, когда все слагаемые не меняются, т.е. все центры кластеров остаются прежними. Такое бывает только при остановке процесса итераций. Другими словами, при продолжении итераций внутрикластерный разброс $g(\Psi)$ монотонно уменьшается. С другой стороны, число различных значений внутрикластерного разброса $g(\Psi)$ конечно, оно не превышает числа разбиений исходного множества элементов $x_1, x_2, \dots, x_n$, на k кластеров. Следовательно, число возможных итераций конечно, рассматриваемый алгоритм всегда остановится, что и требовалось показать.

Итерационные процедуры применяются в различных методах кластер-анализа. Публикация [24] посвящена проблеме остановки алгоритмов — доказательству того, что итерации эталонных алгоритмов (типа «Форель» и метода k -средних) прекращаются через конечное число шагов (оцененное сверху в этой работе). Обобщение было получено в докладе [25]. Итоги многолетних работ по различным вопросам теории классификации подведены в работе [26].

Основные результаты по теории классификации отражены в обширных статьях [15, 27, 28]. Подчеркнем, что все методы классификации, основанные на использовании расстояний (мер различия или близости), естественно рассматривать как часть статистики объектов нечисловой природы [5]. Недавно полученные результаты по теории классификации представлены в работах [29, 30].

ЛИТЕРАТУРА

- Орлов А. И. О развитии математических методов теории классификации / Заводская лаборатория. Диагностика материалов. 2009. Т. 75. № 7. С. 51 – 63.
- Новиков Д. А., Орлов А. И. Математические методы классификации / Заводская лаборатория. Диагностика материалов. 2012. Т. 78. № 4. С. 3.
- Орлов А. И. Математические методы теории классификации / Политематический сетевой электронный научный журнал Кубанского государственного аграрного университета. 2014. № 95. С. 423 – 459.
- Орлов А. И. Структура непараметрической статистики / Заводская лаборатория. Диагностика материалов. 2015. Т. 81. № 7. С. 62 – 72.
- Орлов А. И. Тридцать лет статистики объектов нечисловой природы (обзор) / Заводская лаборатория. Диагностика материалов. 2009. Т. 75. № 5. С. 55 – 64.
- Орлов А. И. Новая парадигма прикладной статистики / Заводская лаборатория. Диагностика материалов. 2012. Т. 78. № 1. Ч. I. С. 87 – 93.
- Леман Э. Л. Проверка статистических гипотез. 2-е изд., испр. — М.: Наука, 1979. — 408 с.
- Орлов А. И. Оценки плотности в пространствах произвольной природы / Статистические методы оценивания и проверки гипотез: межвуз. сб. науч. тр. — Пермь: Перм. гос. нац. иссл. ун-т, 2013. Вып. 25. С. 21 – 33.
- Орлов А. И. Предельные теоремы для ядерных оценок плотности в пространствах произвольной природы / Политематический сетевой электронный научный журнал Кубанского государственного аграрного университета. 2015. № 108. С. 316 – 333.
- Андерсон Т. Введение в многомерный статистический анализ. — М.: Физматгиз, 1963. — 500 с.
- Рао С. Р. Линейные статистические методы и их применения. — М.: Наука, 1968. — 548 с.

12. Алгоритмы и программы восстановления зависимостей / Под ред. В. Я. Вапника. — М.: Наука, 1984. — 816 с.
13. Горелик А. Л., Скрипкин В. А. Методы распознавания: учеб. для вузов. — М.: Высшая школа, 1984. — 208 с.
14. Орлов А. И. Ядерные оценки плотности в пространствах произвольной природы / Статистические методы оценивания и проверки гипотез: межвуз. сб. науч. трудов. — Пермь: Пермский госуниверситет, 1996. С. 68 – 75.
15. Орлов А. И. Математические методы исследования и диагностика материалов / Заводская лаборатория. Диагностика материалов. 2003. Т. 69. № 3. С. 53 – 64.
16. Толчеев В. О. Модифицированный и обобщенный метод ближайшего соседа для классификации библиографических текстовых документов / Заводская лаборатория. Диагностика материалов. 2009. Т. 75. № 7. С. 63 – 70.
17. Алексеевская М. А., Гельфанд И. М., Губерман Ш. А., Мартынов И. В., Ротвайн И. М., Саблин В. М. Прогнозирование исхода мелкоочагового инфаркта миокарда с помощью программы узнавания / Кардиология. 1977. Т. 17. № 7. С. 26 – 71.
18. Гельфанд И. М., Губерман Ш. А., Сыркин А. Л., Головная Л. Д., Извекова М. Л., Алексеевская М. А. Прогнозирование исхода инфаркта миокарда с помощью программы «Кора-3» / Кардиология. 1977. Т. 17. № 6. С. 19 – 23.
19. Гельфанд И. М., Розенфельд Б. И., Шифрин М. А. Очерки о совместной работе математиков и врачей (2-е, дополненное издание). — М.: УРСС, 2004. — 320 с.
20. Фишер Р. Э. Использование множественных измерений в задачах таксономии / Современные проблемы кибернетики. — М.: Знание, 1979. С. 6 – 20.
21. Орлов А. И. Прогностическая сила как показатель качества алгоритма диагностики / Статистические методы оценивания и проверки гипотез: межвуз. сб. науч. тр. — Пермь: Перм. гос. нац. иссл. ун-т, 2011. Вып. 23. С. 104 – 116.
22. Орлов А. И. Прогностическая сила — наилучший показатель качества алгоритма диагностики / Политематический сетевой электронный научный журнал Кубанского государственного аграрного университета. 2014. № 99. С. 15 – 32.
23. Орлов А. И. Организационно-экономическое моделирование: учебник. В 3 ч. Ч. 1. Нечисловая статистика. — М.: Изд-во МГТУ им. Н. Э. Баумана, 2009. — 541 с.
24. Орлов А. И. Сходимость эталонных алгоритмов / Прикладной многомерный статистический анализ: Ученые записки по статистике. — М.: Наука, 1978. Т. 33. С. 361 – 364.
25. Орлов А. И. Остановка после конечного числа шагов для алгоритмов кластер-анализа / Алгоритмическое и программное обеспечение прикладного статистического анализа: Ученые записки по статистике. — М.: Наука, 1980. Т. 36. С. 374 – 377.
26. Орлов А. И. Некоторые вероятностные вопросы теории классификации / Прикладная статистика: Ученые записки по статистике. 1983. Т. 45. С. 166 – 179.
27. Орлов А. И. Классификация объектов нечисловой природы на основе непараметрических оценок плотности / Проблемы компьютерного анализа данных и моделирования: Сборник научных статей. — Минск: Изд-во Белорусского государственного университета, 1991. С. 141 – 148.
28. Орлов А. И. Заметки по теории классификации / Социология: методология, методы, математические модели. 1991. № 2. С. 28 – 50.
29. Орлов А. И., Толчеев В. О. Об использовании непараметрических статистических критериев для оценки точности методов классификации (обобщающая статья) / Заводская лаборатория. Диагностика материалов. 2011. Т. 77. № 3. С. 58 – 66.
30. Орлов А. И. Устойчивость классификации относительно выбора метода кластер-анализа / Заводская лаборатория. Диагностика материалов. 2013. Т. 79. № 1. С. 68 – 71.
4. Orlov A. I. Struktura neparametricheskoi statistiki [Structure of non-parametric statistics] / Zavod. Lab. Diagn. Mater. 2015. Vol. 81. N 7. P. 62 – 72 [in Russian].
5. Orlov A. I. Tridsat' let statistiki ob'ektov nechislvoi prirody (obzor) [Thirty years of statistics of objects of non-numeric nature (review)] / Zavod. Lab. Diagn. Mater. 2009. Vol. 75. N 5. P. 55 – 64 [in Russian].
6. Orlov A. I. Novaya paradigma prikladnoi statistiki [The new paradigm of applied Statistics] / Zavod. Lab. Diagn. Mater. 2012. Vol. 78. N 1. Part I. P. 87 – 93 [in Russian].
7. Leman É. L. Proverka statisticheskikh gipotez [Testing Statistical Hypotheses]. 2nd edition. — Moscow: Nauka, 1979. — 408 p. [in Russian].
8. Orlov A. I. Otsenki plotnosti v prostranstvakh proizvol'noi prirody [Density estimates in spaces of arbitrary nature] / Statisticheskie metody otsenivaniya i proverki gipotez [Statistical methods of estimation and hypothesis testing]: Interschool coll. of sci. papers. — Perm': Izd. Perm. Gos. Nats. Issl. Univ., 2013. Vyp. 25. P. 21 – 33 [in Russian].
9. Orlov A. I. Predel'nye teoremy dlya yadernykh otsenok plotnosti v prostranstvakh proizvol'noi prirody [Limit theorems for kernel density estimators in spaces of arbitrary nature] / Politem. Set. Élektr. Nauch. Zh. Kuban. Gos. Agrar. Univ. 2015. N 108. P. 316 – 333 [in Russian].
10. Anderson T. Vvedenie v mnogomernyi statisticheskii analiz [Introduction to multivariate statistical analysis]. — Moscow: Fizmatgiz, 1963. — 500 p. [in Russian].
11. Rao S. R. Lineinye statisticheskie metody i ikh primeneniya [Linear statistical inference and its applications]. — Moscow: Nauka, 1968. — 548 p. [Russian translation].
12. Vapnik V. Ya. (ed.). Algoritmy i programmy vosstanovleniya zavisimosti [Algorithms and programs for recovery of dependencies]. — Moscow: Nauka, 1984. — 816 p. [in Russian].
13. Gorelik A. L., Skripkin V. A. Metody raspoznavaniya [Methods of recognition]: a textbook for high schools. — Moscow: Vysshaya shkola, 1984. — 208 p. [in Russian].
14. Orlov A. I. Yadernye otsenki plotnosti v prostranstvakh proizvol'noi prirody [The kernel density estimates in spaces of arbitrary nature] / Statisticheskie metody otsenivaniya i proverki gipotez [Statistical methods of estimation and hypothesis testing]. Interschool coll. of sci. papers. — Perm': Izd. PGU, 1996. Vyp. 11. P. 68 – 75 [in Russian].
15. Orlov A. I. Matematicheskie metody issledovaniya i diagnostika materialov [Mathematical methods of research and diagnosis materials (generalizing article)] / Zavod. Lab. Diagn. Mater. 2003. Vol. 69. N 3. P. 53 – 64 [in Russian].
16. Tolcheev V. O. Modifitsirovannyi i obobshchennyi metod blizhaishego sosedya dlya klassifikatsii bibliograficheskikh tekstovykh dokumentov [Modified and generalized method of the nearest neighbor to classify bibliographic text documents] / Zavod. Lab. Diagn. Mater. 2009. Vol. 75. N 7. P. 63 – 70 [in Russian].
17. Alekseevskaya M. A., Gel'fand I. M., Guberman Sh. A., Martynov I. V., Rotvain I. M., Sablin V. M. Prognozirovaniye iskhoda melko-ochagovogo infarkta miokarda s pomoshch'yu programmy uznvaniya [Predicting the outcome of myocardial infarction using recognition program] / Kardiologiya. 1977. Vol. 17. N 7. P. 26 – 71 [in Russian].
18. Gel'fand I. M., Guberman Sh. A., Syркиn A. L., Golovnya L. D., Izvekova M. L., Alekseevskaya M. A. Prognozirovaniye iskhoda infarkta miokarda s pomoshch'yu programmy "Kora-3" [Predicting the outcome of myocardial infarction with the help of the recognition program "Kora-3"] / Kardiologiya. 1977. Vol. 17. N 6. P. 19 – 23 [in Russian].
19. Gelfand I. M., Rosenfeld B. I., Shifrin M. A. Ocherki o sovmestnoi rabote matematikov i vrachei [Essays on joint work of mathematicians and medics]. 2nd edition. — Moscow: URSS, 2004. — 320 p. [Russian translation].
20. Fisher R. A. The use of multiple measurements in taxonomic problems / Ann. Eugenics. 1936. September. Vol. 7. P. 179 – 188.
21. Orlov A. I. Prognosticheskaya sila kak pokazatel' kachestva algoritma diagnostiki [Prognostic strength as an quality indicator of diagnostic algorithm] / Statisticheskie metody otsenivaniya i proverki gipotez [Statistical methods of estimation and hypothesis testing]. Interschool coll. of sci. papers. Vyp. 23. — Perm': Izd. Perm. Gos. Nats. Issl. Univ., 2011. P. 104 – 116 [in Russian].
22. Orlov A. I. Prognosticheskaya sila — nailuchshii pokazatel' kachestva algoritma diagnostiki [Predictive power — the best indicator of the quality of the diagnostic algorithm] / Politem. Set. Élektr. Nauch. Zh. Kuban. Gos. Agrar. Univ. 2014. N 99. P. 15 – 32 [in Russian].
23. Orlov A. I. Organizacionno-ekonomicheskoe modelirovaniye [Organizational-economic modeling]: tutorial. In 3 parts. Part 1. Nchislavaya statistika [non-numerical statistics]. — Moscow: Izd-vo MGTU im. N. É. Bauman, 2009. — 541 p. [in Russian].

REFERENCES

1. Orlov A. I. O razvitiy matematicheskikh metodov teorii klassifikatsii [On the Development of Mathematical Methods in the Theory of Classification (review)] / Zavod. Lab. Diagn. Mater. 2009. Vol. 75. N 7. P. 51 – 63 [in Russian].
2. Novikov D. A., Orlov A. I. Matematicheskie metody klassifikatsii [Mathematical methods of the classification] / Zavod. Lab. Diagn. Mater. 2012. Vol. 78. N 4. P. 3 – 3 [in Russian].
3. Orlov A. I. Matematicheskie metody teorii klassifikatsii [Mathematical methods of the theory of classification] / Politem. Set. Élektr. Nauch. Zh. Kuban. Gos. Agrar. Univ. 2014. N 95. P. 423 – 459 [in Russian].

24. **Orlov A. I.** Skhodimost' étalonnykh algoritmov [The convergence of the reference algorithms] / Prikl. Mnogomer. Stat. Analiz. Uch. Zap. Stat. Vol. 33. — Moscow: Nauka, 1978. P. 361 – 364 [in Russian].
25. **Orlov A. I.** Ostanovka posle konechnogo chisla shagov dlya algoritmov klaster-analiza [Stopping after a finite number of steps for the cluster analysis algorithms] / Algoritm. Progr. Obespech. Prikl. Stat. Analiza. Uch. Zap. Stat. Vol. 36. — Moscow: Nauka, 1980. P. 374 – 377 [in Russian].
26. **Orlov A. I.** Nekotorye veroyatnostnye voprosy teorii klassifikatsii [Some probabilistic questions of the classification theory] / Prikladnaya statistika. Uch. Zap. Stat. Stat. Vol. 45. — Moscow: Nauka, 1983. P. 166 – 179 [in Russian].
27. **Orlov A. I.** Klassifikatsiya ob"ektov nechislovoi prirody na osnove neparametricheskikh otsenok plotnosti [The classification of the objects of non-numeric nature on the basis of nonparametric density estimates] / Problemy komp'yuternogo analiza dannykh i modelirovaniya [Problems of computer data analysis and simulation]: Coll. Sci. Papers. — Minsk: Izd-vo BGU, 1991. P. 141 – 148 [in Russian].
28. **Orlov A. I.** Zametki po teorii klassifikatsii [Notes on the theory of classification] / Sotsiol. Metodol. Metody Matem. Modeli. 1991. N 2. P. 28 – 50 [in Russian].
29. **Orlov A. I., Tolcheev V. O.** Ob ispol'zovanii neparametricheskikh statisticheskikh kriteriev dlya otsenki tochnosti metodov klassifikatsii (obobshchayushchaya stat'ya) [On the use of non-parametric statistical tests to estimate the accuracy of classification methods (generalizing article)] / Zavod. Lab. Diagn. Mater. 2011. Vol. 77. N 3. P. 58 – 66 [in Russian].
30. **Orlov A. I.** Ustoichivost' klassifikatsii otnositel'no vybora metoda klaster-analiza / Zavod. Lab. Diagn. Mater. 2013. Vol. 79. N 1. P. 68 – 71 [in Russian].