

Классификатор дерева решений

Классификатор дерева решений

- Пусть обучающий набор $\mathbf{D} = \{\mathbf{x}_i, y_i\}_{i=1}^n$ состоит из n точек в d -мерном пространстве, причем y_i - метка класса для точки $\mathbf{x}_i$. Мы предполагаем, что измерения или атрибуты X_j являются числовыми или категориальными и что существует k различных классов, так что $y_i \in \{c_1, c_2, \dots, c_k\}$.
- Классификатор дерева решений - это рекурсивная модель дерева на основе секционирования, которая предсказывает класс $\hat{y}_i$ для каждой точки $\mathbf{x}_i$.

- Пусть R обозначает пространство данных, охватывающее множество входных точек D . Дерево решений использует осепараллельную гиперплоскость для разделения пространства данных R на два результирующих полупространства или области, скажем R_1 и R_2 , что также вызывает разбиение входных точек на D_1 и D_2 соответственно. Каждая из этих областей рекурсивно разбивается с помощью осепараллельных гиперплоскостей до тех пор, пока точки внутри индуцированного разбиения не станут относительно чистыми с точки зрения их меток классов, то есть большинство точек принадлежат к одному и тому же классу. Результирующая иерархия разделенных решений образует **модель дерева решений**, с конечными узлами, помеченными классом большинства среди точек в этих регионах.
- Чтобы классифицировать новую тестовую точку, мы должны рекурсивно оценить, к какому полупространству она принадлежит, пока не достигнем листового узла в дереве решений, в котором мы предсказываем ее класс как метку листа

Деревья решений

- **Дерево решений** состоит из внутренних узлов, представляющих решения, соответствующие гиперплоскостям или точкам разделения и конечных узлов, представляющих области или разделы пространства данных, которые помечены классом большинства. Область характеризуется подмножеством точек данных, лежащих в этой области.

Осепараллельные гиперплоскости

- Гиперплоскость $h(\mathbf{x})$ определяется как множество всех точек $\mathbf{x}$, удовлетворяющих следующему уравнению: $h(\mathbf{x}) : \mathbf{w}^T \mathbf{x} + b = 0$ (19.1)
- Здесь $\mathbf{w} \in \mathbb{R}^d$ - **весовой вектор**, нормальный к гиперплоскости, а b - смещение гиперплоскости от начала координат. Дерево решений рассматривает только **осепараллельные гиперплоскости**, то есть весовой вектор должен быть параллелен одному из исходных измерений или осям X_j . Иными словами, весовой вектор $\mathbf{w}$ **априори** ограничен одним из стандартных базисных векторов $\{\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_d\}$, где $\mathbf{e}_i \in \mathbb{R}^d$ имеет 1 для i -го измерения и 0 для всех других измерений. Если $\mathbf{x} = (x_1, x_2, \dots, x_d)^T$ и предполагая $\mathbf{w} = \mathbf{e}_j$, мы можем переписать уравнение (19.1) как $h(\mathbf{x}) : \mathbf{e}_j^T \mathbf{x} + b = 0$, из чего следует $h(\mathbf{x}) : x_j + b = 0$, где выбор смещения b дает различные гиперплоскости вдоль размерности X_j .

Точки Разделения

- Гиперплоскость определяет точку принятия решения или разделения, поскольку она разбивает пространство данных $\mathbf{R}$ на два полупространства. Все точки $\mathbf{x}$, такие, $h(\mathbf{x}) \leq 0$, находятся на гиперплоскости или на одной стороне гиперплоскости, тогда как все точки, такие, что $h(\mathbf{x}) > 0$, находятся на другой стороне. Точки разделения, связанного с осью параллельной гиперплоскости можно записать как $h(\mathbf{x}) \leq 0$, что означает, что $x_i + b \leq 0$ или $x_i \leq -b$. Поскольку x_i - это некоторое значение из измерения X_j и смещение b может быть выбрано как любое значение, общая форма точки разделения для числового атрибута X_j задается как

$$X_j \leq v,$$

где $v = -b$ - некоторое значение в области атрибута X_j . Таким образом, решение или точка разделения $X_j \leq v$ разбивает пространство входных данных $\mathbf{R}$ на две области $\mathbf{R}_Y$ и $\mathbf{R}_N$, которые обозначают множество *всех возможных точек*, удовлетворяющих решению, и тех, которые не удовлетворяют.

Раздел Данных

- Каждое разбиение $\mathbf{R}$ на $\mathbf{R}_Y$ и $\mathbf{R}_N$ также вызывает двоичное разбиение соответствующих входных точек данных $\mathbf{D}$. То есть точка разбиения вида $X_j \leq v$ вызывает разбиение данных

$$\mathbf{D}_Y = \{x \mid x \in \mathbf{D}, x_j \leq v\}$$

$$\mathbf{D}_N = \{x \mid x \in \mathbf{D}, x_j > v\},$$

где $\mathbf{D}_Y$ - подмножество точек данных, лежащих в области $\mathbf{R}_Y$, а $\mathbf{D}_N$ - подмножество входных точек, лежащих в области $\mathbf{R}_N$.

Чистота

- Чистота области $\mathbf{R}_j$ определяется в терминах смеси классов для точек в соответствующем разделе данных $\mathbf{D}_j$. Формально, чистота — это доля точек с меткой большинства в $\mathbf{D}_j$, т. е.

$$\text{purity}(D_j) = \max_i \left\{ \frac{n_{ji}}{n_j} \right\} \quad (19.2),$$

- где $n_j = |\mathbf{D}_j|$ - общее количество точек данных в области $\mathbf{R}_j$, а n_{ji} - количество точек в $\mathbf{D}_j$ с меткой класса c_i .

Категориальные Атрибуты

- Для категориального атрибута X_j точки разделения или решения имеют вид $X_j \in V$, где $V \subset \text{dom}(X_j)$, а $\text{dom}(X_j)$ обозначает область для X_j . Интуитивно это расщепление можно считать категориальным аналогом гиперплоскости. Это приводит к двум "полупространствам", одна область R_Y состоит из точек x , удовлетворяющих условию $x_j \in V$, а другая область R_N состоит из точек, удовлетворяющих условию $x_j \notin V$.

Правила Принятия Решений

- Одним из преимуществ деревьев решений является то, что они создают модели, которые относительно легко интерпретировать. В частности, дерево может быть прочитано как набор решающих правил, причем антецедент каждого правила включает в себя решения по внутренним узлам вдоль пути к листу, а его следствием является метка листового узла. Далее, поскольку все области не пересекаются и охватывают все пространство, набор правил можно интерпретировать как набор альтернатив или дизъюнкций.

Алгоритм дерева решений

- Псевдокод для построения модели дерева решений показан в алгоритме 19.1. Алгоритм принимает в качестве входных данных обучающий набор данных $\mathbf{D}$ и два параметра η и π , где η - размер листа, а π - порог чистоты листа. Различные точки разделения оцениваются для каждого атрибута в $\mathbf{D}$. числовые решения имеют вид $X_j \leq v$ для некоторого значения v в диапазоне значений атрибута X_j , а категориальные решения имеют вид $X_j \in V$ для некоторого подмножества значений в области X_j . Наилучшая точка разбиения выбирается для разбиения данных на два подмножества, $\mathbf{D}_Y$ и $\mathbf{D}_N$, где $\mathbf{D}_Y$ соответствует всем точкам $x \in \mathbf{D}$, удовлетворяющим решению разбиения, а $\mathbf{D}_N$ соответствует всем точкам, не удовлетворяющим решению разбиения. Затем метод дерева решений рекурсивно вызывается на $\mathbf{D}_Y$ и $\mathbf{D}_N$. Для остановки процесса рекурсивного секционирования можно использовать ряд условий остановки. Самое простое условие основано на размере раздела $\mathbf{D}$. Если число точек n в $\mathbf{D}$ падает ниже заданного пользователем порога размера η , то мы останавливаем процесс разбиения и делаем $\mathbf{D}$ листом. Это условие предотвращает чрезмерную подгонку модели к обучающему набору, избегая моделировать очень маленькие подмножества данных. Одного размера недостаточно, потому что, если раздел уже чистый, тогда нет смысла разбивать его дальше. Таким образом, рекурсивное разбиение также прекращается, если чистота D выше порога чистоты π . далее приводится подробная информация о том, как оцениваются и выбираются точки разделения.

ALGORITHM 19.1. Decision Tree Algorithm

DECISIONTREE ($\mathbf{D}, \eta, \pi$):

- 1 $n \leftarrow |\mathbf{D}|$ // partition size
- 2 $n_i \leftarrow |\{\mathbf{x}_j | \mathbf{x}_j \in \mathbf{D}, y_j = c_i\}|$ // size of class c_i
- 3 $\text{purity}(\mathbf{D}) \leftarrow \max_i \left\{ \frac{n_i}{n} \right\}$
- 4 **if** $n \leq \eta$ **or** $\text{purity}(\mathbf{D}) \geq \pi$ **then** // stopping condition
 - 5 $c^* \leftarrow \arg \max_{c_i} \left\{ \frac{n_i}{n} \right\}$ // majority class
 - 6 create leaf node, and label it with class c^*
 - 7 **return**
- 8 $(\text{split point}^*, \text{score}^*) \leftarrow (\emptyset, 0)$ // initialize best split point
- 9 **foreach** (attribute X_j) **do**
 - 10 **if** (X_j is numeric) **then**
 - 11 $(v, \text{score}) \leftarrow \text{EVALUATE-NUMERIC-ATTRIBUTE}(\mathbf{D}, X_j)$
 - 12 **if** $\text{score} > \text{score}^*$ **then** $(\text{split point}^*, \text{score}^*) \leftarrow (X_j \leq v, \text{score})$
 - 13 **else if** (X_j is categorical) **then**
 - 14 $(V, \text{score}) \leftarrow \text{EVALUATE-CATEGORICAL-ATTRIBUTE}(\mathbf{D}, X_j)$
 - 15 **if** $\text{score} > \text{score}^*$ **then** $(\text{split point}^*, \text{score}^*) \leftarrow (X_j \in V, \text{score})$
- // partition $\mathbf{D}$ into $\mathbf{D}_Y$ and $\mathbf{D}_N$ using split point^* , and call recursively
- 16 $\mathbf{D}_Y \leftarrow \{\mathbf{x} \in \mathbf{D} \mid \mathbf{x} \text{ satisfies } \text{split point}^*\}$
- 17 $\mathbf{D}_N \leftarrow \{\mathbf{x} \in \mathbf{D} \mid \mathbf{x} \text{ does not satisfy } \text{split point}^*\}$
- 18 create internal node split point^* , with two child nodes, $\mathbf{D}_Y$ and $\mathbf{D}_N$
- 19 **DECISIONTREE**($\mathbf{D}_Y$); **DECISIONTREE**($\mathbf{D}_N$)

Меры Оценки С Разделением Точек

- Учитывая точку разделения вида $X_j \leq v$ или $X_j \in V$ для числового или категориального атрибута соответственно, нам нужен объективный критерий оценки точки разделения. Интуитивно мы хотим выбрать точку разделения, которая дает лучшее разделение или различие между различными метками классов.

Энтропия

- Энтропия, как правило, измеряет количество беспорядка или неопределенности в системе. В классификационной установке разбиение имеет более низкую энтропию (или низкий беспорядок), если оно относительно чисто, то есть если большинство точек имеют одну и ту же метку. С другой стороны, разделение имеет более высокую энтропию (или больше беспорядка), если метки классов смешаны, и нет класса большинства как такового.

- Энтропия множества помеченных точек $\mathbf{D}$ определяется следующим образом:

$$H(D) = - \sum_{i=1}^k P(C_i | \mathbf{D}) \log_2 P(c_i | \mathbf{D}) \quad (19.3),$$

- где $P(c_i | \mathbf{D})$ - вероятность класса c_i в $\mathbf{D}$, А k -число классов. Если область чиста, то есть имеет точки из одного класса, то энтропия равна нулю. С другой стороны, если все классы перемешаны и каждый появляется с равной вероятностью $P(c_i | \mathbf{D}) = \frac{1}{k}$, то энтропия имеет наибольшее значение, $H(\mathbf{D}) = \log_2 k$.
- Предположим, что точка разделения разбивает $\mathbf{D}$ на $\mathbf{D}_Y$ и $\mathbf{D}_N$. Определите энтропию разделения как взвешенную энтропию каждого из результирующих разбиений, заданную как

$$H(\mathbf{D}_Y, \mathbf{D}_N) = \frac{n_Y}{n} H(\mathbf{D}_Y) + \frac{n_N}{n} H(\mathbf{D}_N) \quad (19.4),$$

- где $n = |\mathbf{D}|$ - количество точек в $\mathbf{D}$, а $n_Y = |\mathbf{D}_Y|$ и $n_N = |\mathbf{D}_N|$ - количество точек в $\mathbf{D}_Y$ и $\mathbf{D}_N$.

- Чтобы увидеть, приводит ли точка разделения к уменьшению общей энтропии, мы определяем информационный прирост для данной точки разделения следующим образом:

$$Gain(\mathbf{D}, \mathbf{D}_Y, \mathbf{D}_N) = H(\mathbf{D}) - H(\mathbf{D}_Y, \mathbf{D}_N) \quad (19.5)$$

- Чем выше прирост информации, тем больше уменьшение энтропии и тем лучше точка разделения. Таким образом, учитывая точки разбиения и соответствующие им разбиения, мы можем оценить каждую точку разбиения и выбрать ту, которая дает наибольший информационный выигрыш.

Индекс Джини

- Другой распространенной мерой для оценки чистоты точки разделения является индекс Джини, определяемый следующим образом:

$$G(\mathbf{D}) = 1 - \sum_{i=1}^k P(c_i|\mathbf{D})^2 \quad (19.6)$$

- Если разбиение чисто, то вероятность мажоритарного класса равна 1, а вероятность всех остальных классов равна 0, и, таким образом, индекс Джини равен 0. С другой стороны, когда каждый класс представлен одинаково, с вероятностью $P(c_i|\mathbf{D}) = \frac{1}{k}$, то индекс Джини имеет значение $\frac{k-1}{k}$. Таким образом, более высокие значения индекса Джини указывают на больший беспорядок, а более низкие значения указывают на больший порядок с точки зрения меток классов.

- Мы можем вычислить взвешенный индекс Джини точки разделения следующим образом:

$$G(\mathbf{D}_Y, \mathbf{D}_N) = \frac{n_Y}{n} G(\mathbf{D}_Y) + \frac{n_N}{n} G(\mathbf{D}_N),$$

- где n , n_Y и n_N обозначают количество точек в областях $\mathbf{D}$, $\mathbf{D}_Y$ и $\mathbf{D}_N$ соответственно. Чем ниже значение индекса Джини, тем лучше точка разделения.
- Вместо энтропии и индекса Джини для оценки расщеплений можно использовать и другие показатели. Например, мера деревьев классификации и регрессии (CART) задается следующим образом

$$CART(\mathbf{D}_Y, \mathbf{D}_N) = 2 \frac{n_Y}{n} \frac{n_N}{n} \sum_{i=1}^k |P(c_i | \mathbf{D}_Y) - P(c_i | \mathbf{D}_N)| \quad (19.7)$$

- Таким образом, эта мера предпочитает точку разбиения, которая максимизирует разницу между массовой функцией вероятности класса для двух разбиений; чем выше мера CART, тем лучше точка разбиения.

Оценка Точек Разделения

- Все меры оценки точки разделения, такие как энтропия [Ур. (19.3)], индекс Джини [Ур. (19.6)] и CART [Ур. (19.7)], рассмотренные в предыдущем разделе, зависят от класса вероятностной массовой функции (ВМФ) для $\mathbf{D}$, а именно $P(c_i | \mathbf{D})$, и класса ВМФ для результирующих разбиений $\mathbf{D}_Y$ и $\mathbf{D}_N$, а именно $P(c_i | \mathbf{D}_Y)$ и $P(c_i | \mathbf{D}_N)$. Обратите внимание, что мы должны вычислить класс ВМФ для всех возможных точек разделения; оценка каждой из них независимо приведет к значительным вычислительным затратам. Вместо этого можно постепенно вычислить ВМФ, как описано в следующих параграфах.

Числовые Атрибуты

- Если X является числовым атрибутом, мы должны оценить точки разделения вида $X \leq v$. даже если мы ограничим v лежащим в пределах диапазона значений атрибута X , для V все равно существует бесконечное число вариантов выбора. Один из разумных подходов состоит в том, чтобы рассматривать только средние точки между двумя последовательными различными значениями X в выборке D . Это происходит потому, что точки разбиения вида $X \leq v$ для $v \in [x_a, x_b)$, где x_a и x_b - два последовательных различных значения X в D , производят одно и то же разбиение D на D_Y и D_N и, таким образом, дают одни и те же оценки. Поскольку для X может быть не более n различных значений, необходимо учитывать не более $n-1$ средних значений.

- Пусть $\{v_1, \dots, v_m\}$ обозначает множество всех таких средних точек, таких что $v_1 < v_2 < \dots < v_m$. Для каждой точки $x \leq V$, то есть, чтобы оценить класс ВМФ:

$$\hat{P}(c_i | \mathbf{D}_Y) = \hat{P}(c_i | X \leq v) \quad (19.8)$$

$$\hat{P}(c_i | \mathbf{D}_Y) = \hat{P}(c_i | X > v) \quad (19.9)$$

- Пусть $I()$ - индикаторная переменная, которая принимает значение 1 только тогда, когда ее аргумент истинен, и 0 в противном случае. Используя теорему Байеса, мы имеем

$$\hat{P}(c_i | X \leq v) = \frac{\hat{P}(X \leq v | c_i) \hat{P}(c_i)}{\hat{P}(X \leq v)} = \frac{\hat{P}(X \leq v | c_i) \hat{P}(c_i)}{\sum_{j=1}^k \hat{P}(X \leq v | c_j) \hat{P}(c_j)} \quad (19.10)$$

- Априорная вероятность для каждого класса в $\mathbf{D}$ может быть оценена следующим образом:

$$\hat{P}(c_i) = \frac{1}{n} \sum_{j=1}^n I(y_j = c_i) = \frac{n_i}{n} \quad (19.11)$$

- где y_j - класс для точки $\mathbf{x}_j$, $n = |\mathbf{D}|$ - общее количество точек, а n_i - количество точек в $\mathbf{D}$ с классом c_i .

- Определите N_{vi} как число точек $x_j \leq v$ с классом c_i , где x_j - значение точки данных x_j для атрибута X , заданное как

$$N_{vi} = \sum_{j=1}^n I(x_j \leq v \text{ and } y_i = c_i) \quad (19.12)$$

- Затем мы можем оценить $P(X \leq v | c_i)$ следующим образом:

$$\hat{P}(X \leq v | c_i) = \frac{\hat{P}(X \leq v \text{ and } c_i)}{\hat{P}(c_i)} = \left(\frac{1}{n} \sum_{j=1}^n I(x_j \leq v \text{ and } y_j = c_i) \right) / (n_i/n) = \frac{N_{vi}}{n_i} \quad (19.13)$$

- Подставляя уравнения. (19.11) и (19.13) в уравнение (19.10), и используя уравнение (19.8), мы имеем

$$\hat{P}(c_i | \mathbf{D}_Y) = \hat{P}(c_i | X \leq v) = \frac{N_{vi}}{\sum_{j=1}^k N_{vj}} \quad (19.14)$$

- Мы можем оценить $\hat{P}(X > v | c_i)$ следующим образом:

$$\hat{P}(X > v | c_i) = 1 - \hat{P}(X \leq v | c_i) = 1 - \frac{N_{vi}}{n_i} = \frac{n_i - N_{vi}}{n_i} \quad (19.15)$$

- Используя уравнения (19.11), (19.15) класс ВМФ $\hat{P}(c_i | \mathbf{D}_N)$ рассчитывается:

$$\hat{P}(c_i | \mathbf{D}_N) = \hat{P}(c_i | X > v) = \frac{\hat{P}(X > v | c_i) \hat{P}(c_i)}{\sum_{j=1}^k \hat{P}(X > v | c_j) \hat{P}(c_j)} = \frac{n_i - N_{vi}}{\sum_{j=1}^k (n_j - N_{vj})} \quad (19.16)$$

Алгоритм

- Алгоритм 19.2 показывает метод оценки точки разделения для числовых атрибутов. Цикл for в строке 4 перебирает все точки и вычисляет значения средних точек v и количество точек N_{v_i} из класса c_i таким образом, что $x_j \leq v$. Цикл for в строке 12 перечисляет все возможные точки разделения вида $X \leq v$, по одной для каждой средней точки v , и оценивает их с помощью критерия усиления [Ур. (19.5)]; лучшая точка разделения и оценка записываются и возвращаются. Можно также использовать любые другие оценочные меры. Однако для индекса Джини и CART более низкий балл лучше, в отличие от прироста, где более высокий балл лучше.

Сложность

- С точки зрения вычислительной сложности начальная сортировка значений X (строка 1) занимает время $O(n \log n)$. Стоимость вычисления средних точек и специфичных для класса отсчетов N_{vi} занимает время $O(nk)$ (для цикла в строке 4). Стоимость вычисления балла также ограничена $O(nk)$, поскольку общее число средних точек v может быть не более n (для цикла в строке 12). Таким образом, общая стоимость оценки числового атрибута равна $O(n \log n + nk)$. Игнорируя k , поскольку это обычно небольшая константа, общая стоимость вывода числовой оценки точки разделения равна $O(n \log n)$.

ALGORITHM 19.2. Evaluate Numeric Attribute (Using Gain)

EVALUATE-NUMERIC-ATTRIBUTE ($\mathbf{D}, X$):

```
1 sort  $\mathbf{D}$  on attribute  $X$ , so that  $x_j \leq x_{j+1}, \forall j = 1, \dots, n-1$ 
2  $\mathcal{M} \leftarrow \emptyset$  // set of midpoints
3 for  $i = 1, \dots, k$  do  $n_i \leftarrow 0$ 
4 for  $j = 1, \dots, n-1$  do
5   if  $y_j = c_i$  then  $n_i \leftarrow n_i + 1$  // running count for class  $c_i$ 
6   if  $x_{j+1} \neq x_j$  then
7      $v \leftarrow \frac{x_{j+1} + x_j}{2}; \mathcal{M} \leftarrow \mathcal{M} \cup \{v\}$  // midpoints
8     for  $i = 1, \dots, k$  do
9        $N_{vi} \leftarrow n_i$  // Number of points such that  $x_j \leq v$  and  $y_j = c_i$ 
10 if  $y_n = c_i$  then  $n_i \leftarrow n_i + 1$ 
    // evaluate split points of the form  $X \leq v$ 
11  $v^* \leftarrow \emptyset; score^* \leftarrow 0$  // initialize best split point
12 forall  $v \in \mathcal{M}$  do
13   for  $i = 1, \dots, k$  do
14      $\hat{P}(c_i | \mathbf{D}_Y) \leftarrow \frac{N_{vi}}{\sum_{j=1}^k N_{vj}}$ 
15      $\hat{P}(c_i | \mathbf{D}_N) \leftarrow \frac{n_i - N_{vi}}{\sum_{j=1}^k n_j - N_{vj}}$ 
16    $score(X \leq v) \leftarrow Gain(\mathbf{D}, \mathbf{D}_Y, \mathbf{D}_N)$  // use Eq. (19.5)
17   if  $score(X \leq v) > score^*$  then
18      $v^* \leftarrow v; score^* \leftarrow score(X \leq v)$ 
19 return  $(v^*, score^*)$ 
```

Категориальные Атрибуты

- Если X является категориальным атрибутом, то мы оцениваем точки расщепления вида $X \in V$, где $V \subset \text{dom}(X)$ и $V \neq \emptyset$. Другими словами, рассматриваются все различные разбиения множества значений X . Поскольку точка разделения $X \in V$ дает то же разбиение, что и $X \in V$, где $V = \text{dom}(X) \setminus V$ является дополнением к V , общее число различных разбиений задается как

$$\sum_{i=1}^{\lfloor m/2 \rfloor} \binom{m}{i} = O(2^{m-1}) \quad (19.17),$$

- где m -число значений в области X , то есть $m = |\text{dom}(X)|$. Таким образом, число возможных точек расщепления для рассмотрения экспоненциально в m , что может создавать проблемы, если m велико. Одно из упрощений состоит в том, чтобы ограничить V размером один, так что существует только m расщепленных точек вида $X_j \in \{v\}$, где $v \in \text{dom}(X_j)$.

- Чтобы оценить данный раскол точке $x \in V$, мы должны вычислить следующую вероятность массового функций класса:

$$P(c_i | \mathbf{D}_Y) = P(c_i | X \in V)$$

$$P(c_i | \mathbf{D}_N) = P(c_i | X \notin V)$$

- Используя теорему Байеса, мы имеем

$$P(c_i | X \in V) = \frac{P(X \in V | c_i) P(c_i)}{P(X \in V)} = \frac{P(X \in V | c_i) P(c_i)}{\sum_{j=1}^k P(X \in V | c_j) P(c_j)}$$

- Однако заметим, что данная точка x может принимать только одно значение в области X , и поэтому значения $v \in \text{dom}(X)$ являются взаимоисключающими. Поэтому мы имеем

$$P(X \in V | c_i) = \sum_{v \in V} P(X = v | c_i)$$

- и мы можем переписать $P(c_i | \mathbf{D}_Y)$ как

$$P(c_i | \mathbf{D}_Y) = \frac{\sum_{v \in V} P(X=v | c_i) P(c_i)}{\sum_{j=1}^k \sum_{v \in V} P(X=v | c_j) P(c_j)} \quad (19.18)$$

- Определим n_{vi} как число точек $\mathbf{x}_j \in \mathbf{D}$ со значением $x_j = v$ для атрибута X и классом $y_j = c_i$:

$$n_{vi} = \sum_{j=1}^n I(x_j = v \text{ and } y_j = c_i) \quad (19.19)$$

- Класс условных эмпирических ВМФ для X тогда задается как

$$\hat{P}(X = v | c_i) = \frac{\hat{P}(X=v \text{ and } c_i)}{\hat{P}(c_i)} = \left(\frac{1}{n} \sum_{j=1}^n I(x_j = v \text{ and } y_j = c_i) \right) / (c_i / n) = n_{vi} / n_i \quad (19.20)$$

- Заметим, что априорные вероятности класса могут быть оценены с помощью уравнения (19.11), как обсуждалось для разбиения $\mathbf{D}_Y$ для точки разбиения $X \in V$ задается как

$$\hat{P}(c_i | \mathbf{D}_Y) = \frac{\sum_{v \in V} \hat{P}(X=v | c_i) \hat{P}(c_i)}{\sum_{j=1}^k \sum_{v \in V} \hat{P}(X=v | c_j) \hat{P}(c_j)} = \frac{\sum_{v \in V} n_{vi}}{\sum_{j=1}^k \sum_{v \in V} n_{vj}} \quad (19.21)$$

- Аналогичным образом класс ВМФ для раздела $\mathbf{D}_N$ задается как

$$\hat{P}(c_i | \mathbf{D}_N) = \hat{P}(c_i | X \notin V) = \frac{\sum_{v \notin V} n_{vi}}{\sum_{j=1}^k \sum_{v \notin V} n_{vj}} \quad (19.22)$$

Алгоритм

- Алгоритм 19.3 показывает метод оценки точки разделения для категориальных атрибутов. Цикл for в строке 4 перебирает все точки и вычисляет n_{vi} , то есть число точек, имеющих значение $v \in \text{dom}(X)$ и класс c_i . Цикл for в строке 7 перечисляет все возможные точки разбиения вида $X \in V$ для $V \subset \text{dom}(X)$, такие что $|V| \leq l$, где l - заданный пользователем параметр, обозначающий максимальную мощность V . Например, чтобы контролировать количество точек разбиения, мы также можем ограничить V одним элементом, то есть $l = 1$, так что разбиения имеют вид $V \in \{v\}$, с $v \in \text{dom}(X)$. Если $l = \lfloor m/2 \rfloor$, то мы должны рассмотреть все возможные различные разбиения V . учитывая точку разбиения $X \in V$, метод оценивает ее с помощью информационного усиления [Ур. (19.5)], хотя также можно использовать любой другой критерий оценки. Лучшая точка разделения и оценка записываются и возвращаются.

Сложность

- С точки зрения вычислительной сложности специфичные для класса подсчеты для каждого значения n_{v_i} занимают $O(n)$ времени (для цикла в строке 4). При $m = |\text{dom}(X)|$ максимальное число разбиений V равно $O(2^{m-1})$, и поскольку каждая точка разбиения может быть оценена за время $O(mk)$, цикл for В строке 7 занимает время $O(mk2^{m-1})$. Таким образом, общая стоимость категориальных атрибутов равна $O(n + mk2^{m-1})$. Если мы предположим, что $2^{m-1} = O(n)$, то есть если мы связываем максимальный размер V с $I = O(\log n)$, то стоимость категориальных расщеплений ограничена как $O(n \log n)$, игнорируя k .

ALGORITHM 19.3. Evaluate Categorical Attribute (Using Gain)

EVALUATE-CATEGORICAL-ATTRIBUTE ($\mathbf{D}, X, l$):

```
1 for  $i = 1, \dots, k$  do
2    $n_i \leftarrow 0$ 
3   forall  $v \in \text{dom}(X)$  do  $n_{vi} \leftarrow 0$ 
4 for  $j = 1, \dots, n$  do
5   if  $x_j = v$  and  $y_j = c_i$  then  $n_{vi} \leftarrow n_{vi} + 1$  // frequency statistics
   // evaluate split points of the form  $X \in V$ 
6  $V^* \leftarrow \emptyset; \text{score}^* \leftarrow 0$  // initialize best split point
7 forall  $V \subset \text{dom}(X)$ , such that  $1 \leq |V| \leq l$  do
8   for  $i = 1, \dots, k$  do
9      $\hat{P}(c_i | \mathbf{D}_Y) \leftarrow \frac{\sum_{v \in V} n_{vi}}{\sum_{j=1}^k \sum_{v \in V} n_{vj}}$ 
10     $\hat{P}(c_i | \mathbf{D}_N) \leftarrow \frac{\sum_{v \notin V} n_{vi}}{\sum_{j=1}^k \sum_{v \notin V} n_{vj}}$ 
11     $\text{score}(X \in V) \leftarrow \text{Gain}(\mathbf{D}, \mathbf{D}_Y, \mathbf{D}_N)$  // use Eq. (19.5)
12    if  $\text{score}(X \in V) > \text{score}^*$  then
13       $V^* \leftarrow V; \text{score}^* \leftarrow \text{score}(X \in V)$ 
14 return  $(V^*, \text{score}^*)$ 
```

Вычислительная Сложность

- Для анализа вычислительной сложности метода дерева решений в алгоритме 19.1 мы предполагаем, что стоимость оценки всех точек разбиения для числового или категориального атрибута - $O(n \log n)$, где $n = |\mathbf{D}|$ - размер набора данных. Учитывая $\mathbf{D}$, алгоритм дерева решений оценивает все атрибуты d со стоимостью $(d n \log n)$. Общая стоимость зависит от глубины дерева решений. В худшем случае дерево может иметь глубину n , и таким образом общая стоимость равна $O(d n^2 \log n)$.