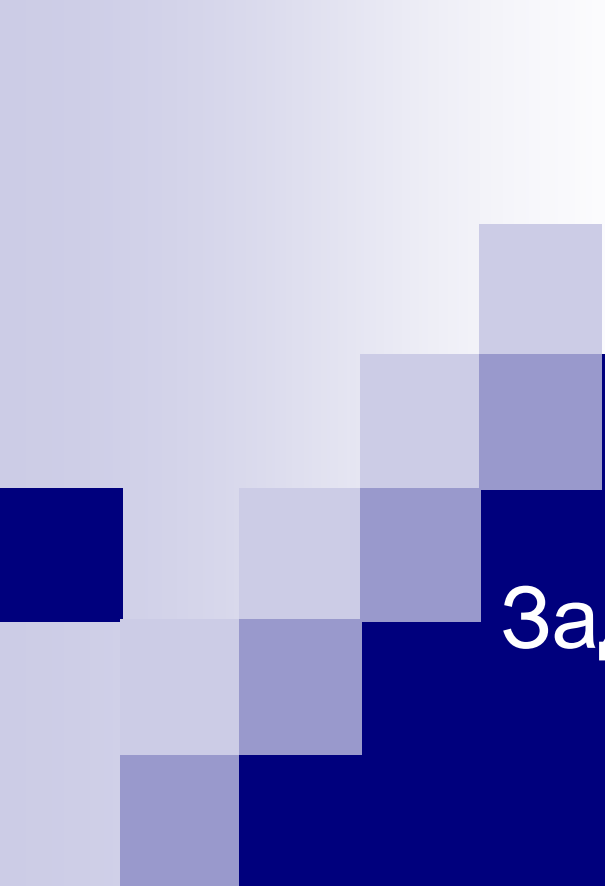




# Лекция 7: Задача классификации и деревья решений

# Задачи классификации

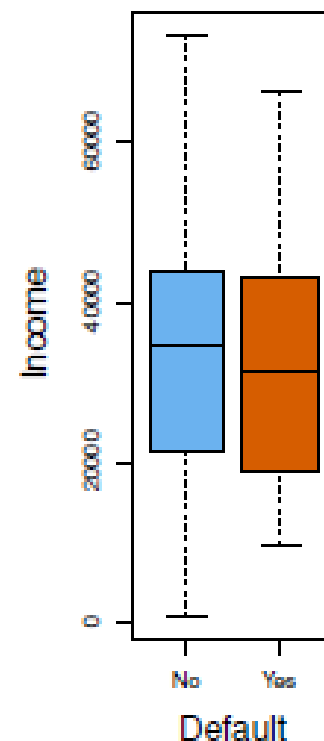
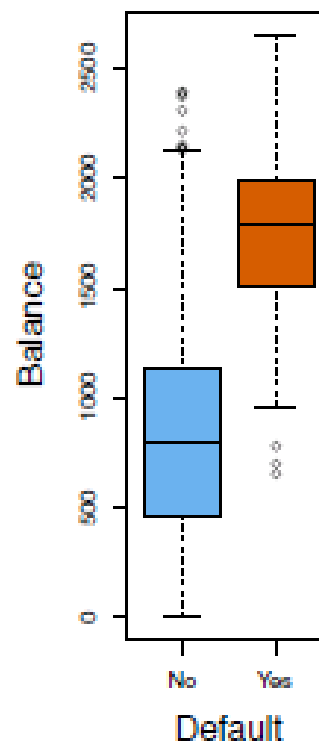
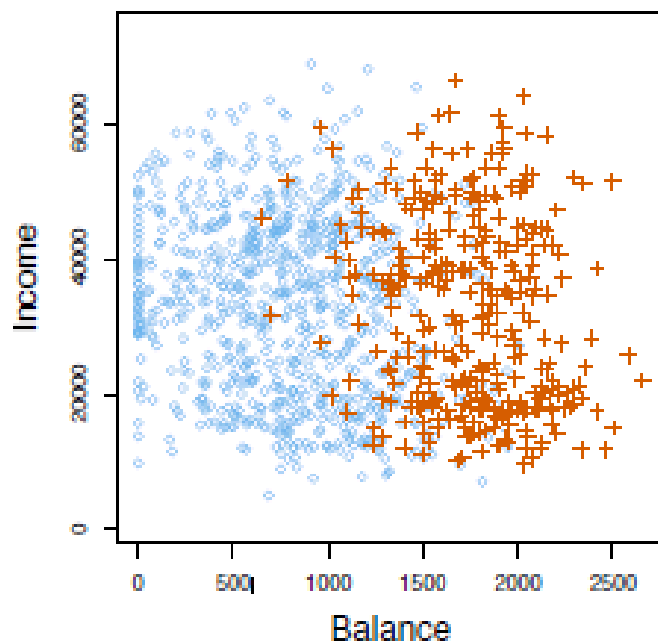
Переменная отклика  $Y$  является *категориальной* – например, почтовое сообщение может принадлежать одному из следующих классов  $C = (\text{spam}; \text{ham})$  ( $\text{ham}$  = обычная почта), изображение цифры может принадлежать одному из классов  $C = \{0, 1, \dots, 9\}$ .

Цель задачи:

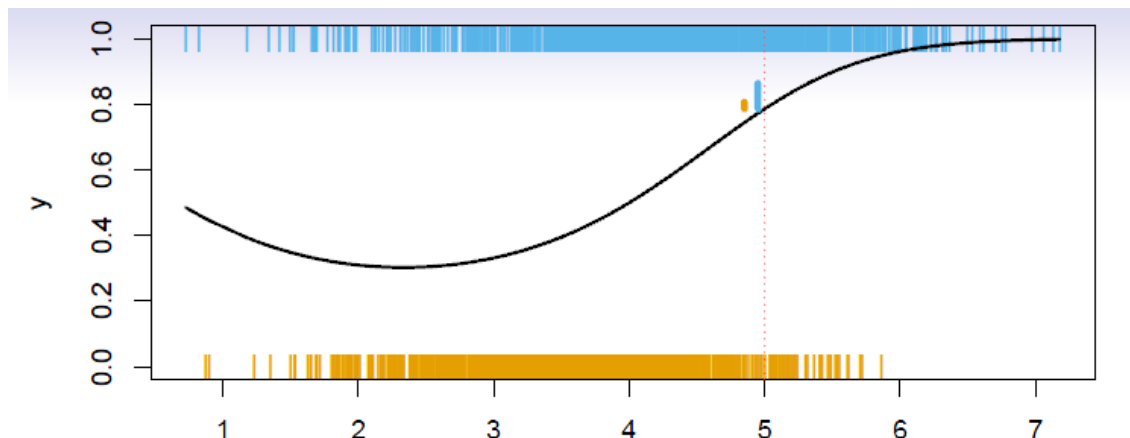
- Построить классификатор  $C(X)$ , который связывает метку класса из множества  $C$  с неразмеченным образцом  $X$ .
- Оценка степени неточности для каждой классификации
- Понимание роли различных предикторов среди

$$X = (X_1, X_2, \dots, X_p).$$

# Пример: Банкротство по кредитным картам



# Задача классификации



Существует ли идеальная  $C(X)$ ? Пусть  $K$  классов в множестве  $S$  пронумерованы как  $1, 2, \dots, K$ . Пусть

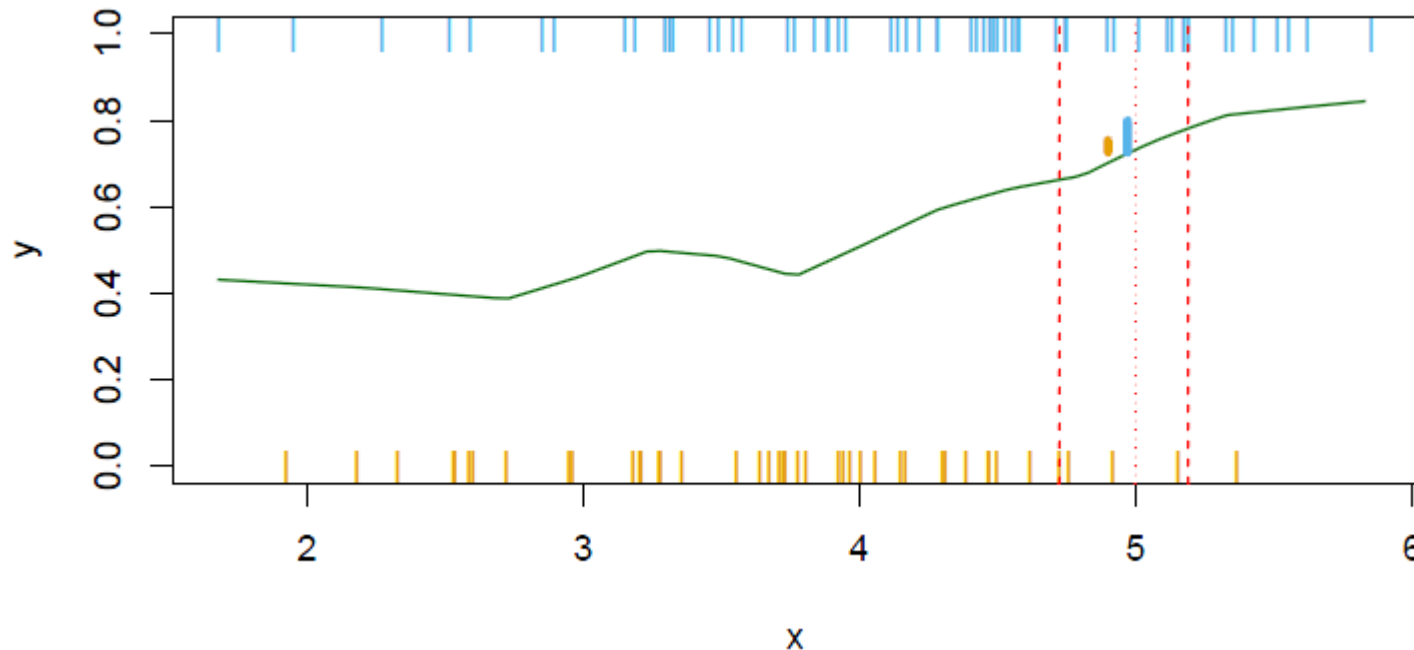
$$p_k(x) = \Pr(Y = k | X = x), \quad k = 1, 2, \dots, K.$$

Это условные вероятности классов в точке  $x$ . Тогда оптимальный классификатор Байеса в точке  $x$ :

$$C(x) = j \text{ if } p_j(x) = \max\{p_1(x), p_2(x), \dots, p_K(x)\}$$

Качество классификации:  $\text{Err}_{T_e} = \text{Ave}_{i \in T_e} I[y_i \neq \hat{C}(x_i)]$

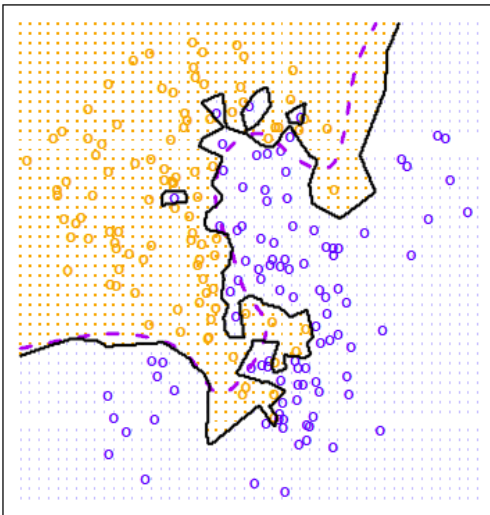
# Задача классификации KNN



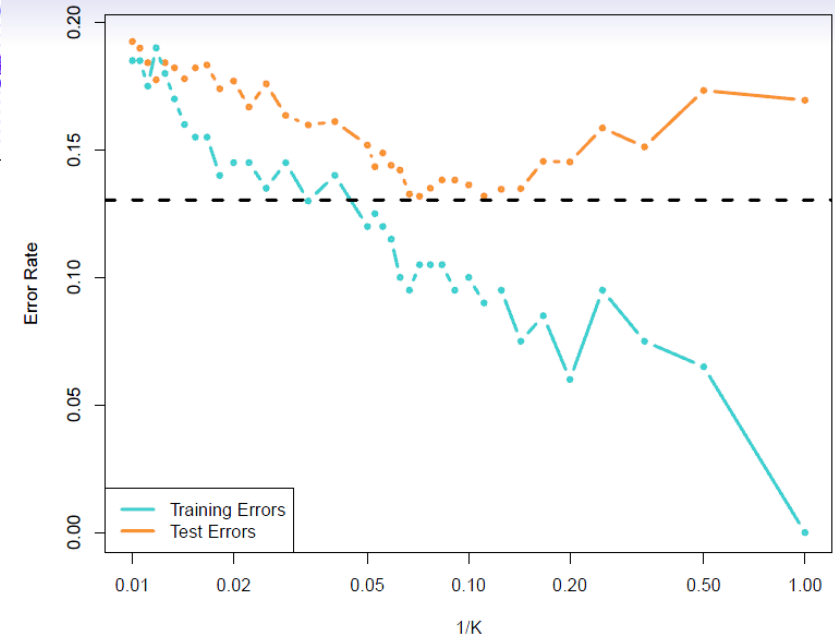
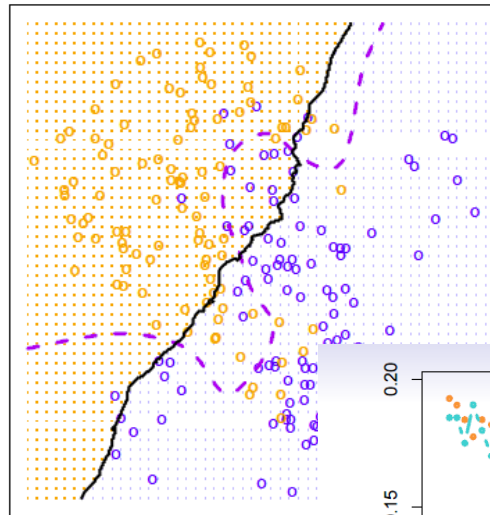
- Также может быть использовано усреднение методом ближайших соседей.
- Хотя такой подход плохо применим при возрастании размерности.

# Пример: KNN для двух измерений

KNN: K=1



KNN: K=100



# Методы Байеса

- Вероятностная постановка задачи прогнозирования:
  - В явном виде вычисляется вероятность для каждой гипотезы (варианта прогноза), очень важно для многих задач
- Возможность дообучения:
  - Каждый размеченный пример итеративно увеличивает/уменьшает вероятность корректности гипотезы
  - Априорные знания легко внедрить в модель
- Вероятностный прогноз:
  - Несколько вариантов гипотез с оценкой достоверности
- Де факто - эталон:
  - Даже когда методы Байеса вычислительно трудоемки, они могут быть полезны в качестве разовой оценки для сравнения с более быстрыми методами в плане оценки точности последних

# Теорема Байеса

- Теореме Байеса:

- Если  $Z$  - тренировочный набор, то апостериорная вероятность гипотезы  $h$ ,  $P(h|Z)$  удовлетворяет соотношению:

$$P(h|Z) = \frac{P(Z|h)P(h)}{P(Z)}$$

- MAP (maximum posteriori) гипотеза

$$h_{MAP} \equiv \arg \max_{h \in H} P(h|Z) = \arg \max_{h \in H} P(Z|h)P(h).$$

- «Необходимые» оценки вероятностей:

- $P(Z)$  – константа для всех классов
  - $P(h)$  - относительная частота класса в выборке
  - Надо найти  $h$  такое, что  $P(h|Z)$  максимально =  $h$  такое, что  $P(Z|h) \cdot P(h)$  максимально

- Проблема:

- вычисление  $P(Z|h)$  на всех комбинациях переменных вычислительно сложно или невозможно

# Метод Naïve Байеса

- Naïve упрощение:

- Атрибуты независимы:  $P(x_1, \dots, x_n | C) = P(x_1 | C) \cdot \dots \cdot P(x_n | C)$

- Оценка  $P(x_i | C)$ :

- Если атрибут категориальный, то относительная частота того, сколько примеров с классом  $C$  имеют значение атрибута  $x_i$
  - Если атрибут числовой, то по плотности распределения, например:

$$f_k(x) = \frac{1}{\sqrt{2\pi}\sigma_k} e^{-\frac{1}{2}\left(\frac{x-\mu_k}{\sigma_k}\right)^2}$$

где  $\mu_k$  – это среднее, и  $\sigma_k^2$  - дисперсия (в классе  $k$ ).

- И так, и так считается быстро

- Предположение о независимости

- Делает метод вычислительно эффективным
  - Приводит к оптимальному классификатору (если и правда независимые атрибуты), но на практике почти всегда зависимы

- Попытки преодолеть эти ограничения:

- Байесовские сети

# Пример: построить модель прогноза ВОЗМОЖНОСТИ ПОИГРАТЬ В ТЕННИС

| Outlook  | Temperature | Humidity | Windy | Class |
|----------|-------------|----------|-------|-------|
| sunny    | hot         | high     | false | N     |
| sunny    | hot         | high     | true  | N     |
| overcast | hot         | high     | false | P     |
| rain     | mild        | high     | false | P     |
| rain     | cool        | normal   | false | P     |
| rain     | cool        | normal   | true  | N     |
| overcast | cool        | normal   | true  | P     |
| sunny    | mild        | high     | false | N     |
| sunny    | cool        | normal   | false | P     |
| rain     | mild        | normal   | false | P     |
| sunny    | mild        | normal   | true  | P     |
| overcast | mild        | high     | true  | P     |
| overcast | hot         | normal   | false | P     |
| rain     | mild        | high     | true  | N     |

$$P(p) = 9/14$$

$$P(n) = 5/14$$

| outlook                      |                            |
|------------------------------|----------------------------|
| $P(\text{sunny} p) = 2/9$    | $P(\text{sunny} n) = 3/5$  |
| $P(\text{overcast} p) = 4/9$ | $P(\text{overcast} n) = 0$ |
| $P(\text{rain} p) = 3/9$     | $P(\text{rain} n) = 2/5$   |
| temperature                  |                            |
| $P(\text{hot} p) = 2/9$      | $P(\text{hot} n) = 2/5$    |
| $P(\text{mild} p) = 4/9$     | $P(\text{mild} n) = 2/5$   |
| $P(\text{cool} p) = 3/9$     | $P(\text{cool} n) = 1/5$   |
| humidity                     |                            |
| $P(\text{high} p) = 3/9$     | $P(\text{high} n) = 4/5$   |
| $P(\text{normal} p) = 6/9$   | $P(\text{normal} n) = 2/5$ |
| windy                        |                            |
| $P(\text{true} p) = 3/9$     | $P(\text{true} n) = 3/5$   |
| $P(\text{false} p) = 6/9$    | $P(\text{false} n) = 2/5$  |

# Пример

- Новый пример:

- $x = \langle \text{rain, hot, high, false} \rangle$

- Расчет вероятностей:

- $P(x|p) \cdot P(p) = P(\text{rain}|p) \cdot P(\text{hot}|p) \cdot P(\text{high}|p) \cdot P(\text{false}|p) \cdot P(p) = 3/9 \cdot 2/9 \cdot 3/9 \cdot 6/9 \cdot 9/14 = 0.010582$

- $P(x|n) \cdot P(n) = P(\text{rain}|n) \cdot P(\text{hot}|n) \cdot P(\text{high}|n) \cdot P(\text{false}|n) \cdot P(n) = 2/5 \cdot 2/5 \cdot 4/5 \cdot 2/5 \cdot 5/14 = 0.018286$

- Вывод:

- $P(x|n) \cdot P(n) > P(x|p) \cdot P(p)$

- Значит скорее всего не поиграть

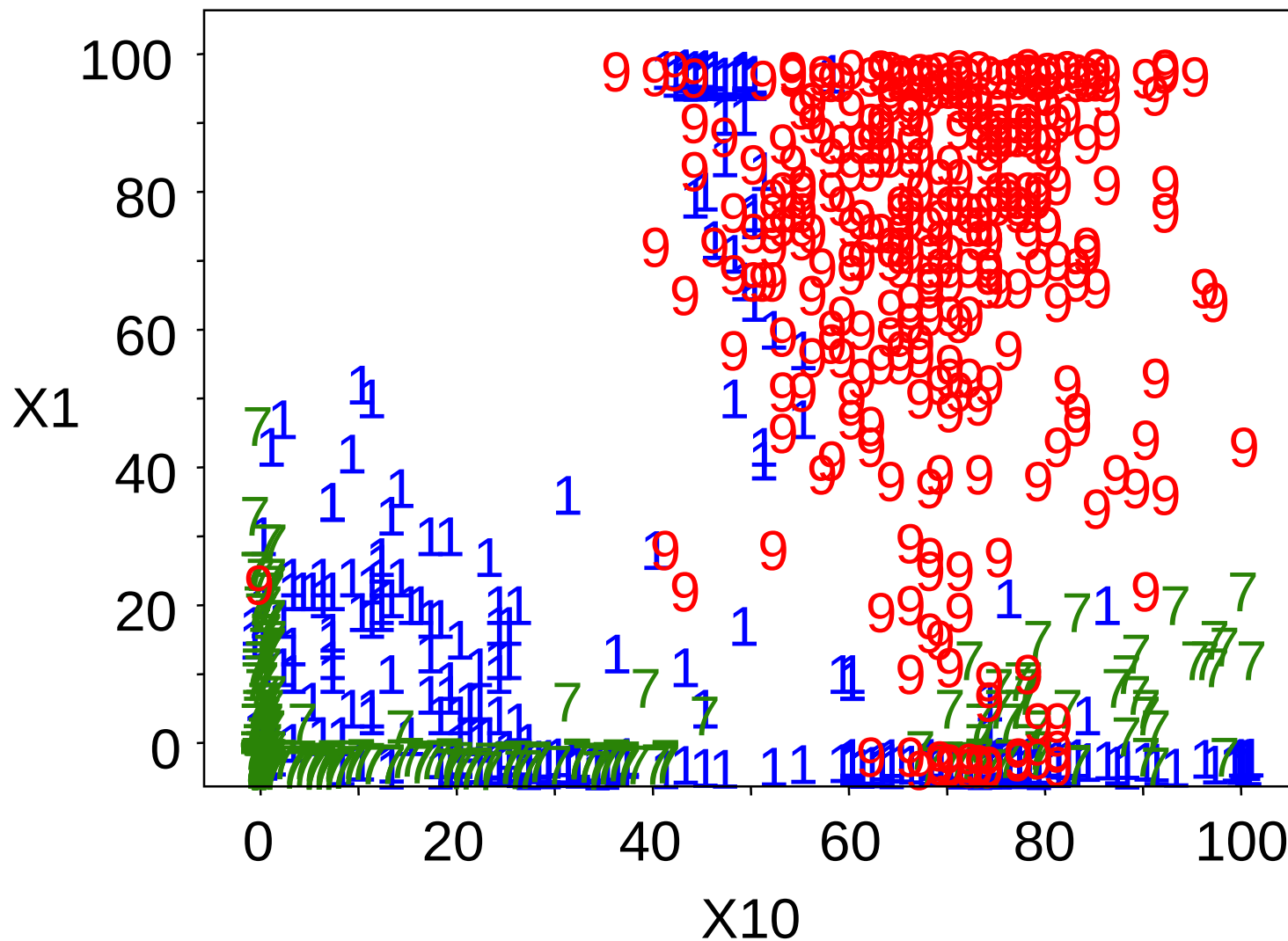
# Методы, основанные на деревьях решений

- Эти методы используют *стратификацию или сегментирование* пространства признаков на области.
- Для сегментации пространства признаков может использоваться набор правил, который можно представить в виде дерева
- Деревья решений могут применяться как к задачам регрессии, так и к классификации.
- Методы, основанные на деревьях, просты в интерпретации, при этом показывают достаточно хорошие результаты по точности прогнозирования.
- Нестабильные модели – это плюс для ансамблей *бэггинг, методы случайного леса и бустинг*. Эти методы строят множество деревьев, результаты прогнозирования которых потом объединяются для получения итогового прогноза.

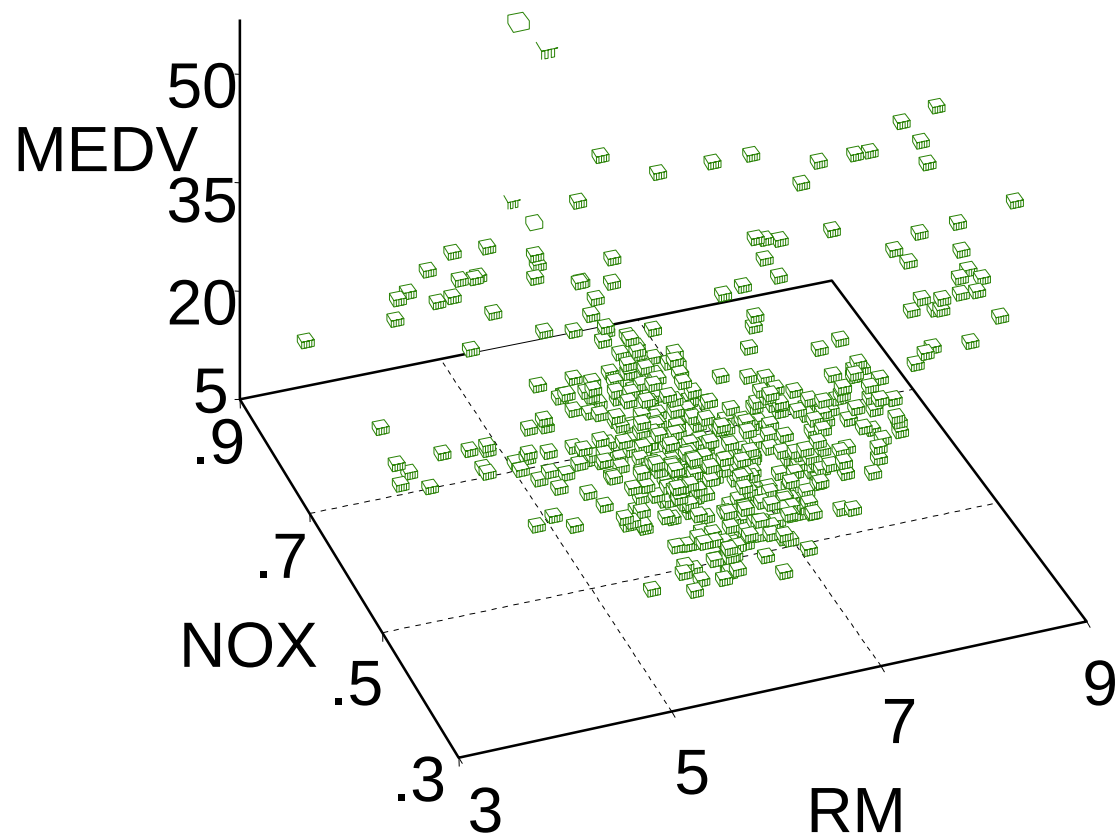
# Деревья решений в задачах классификации и регрессии

- Дерево решений - граф (древовидная структура), в котором:
  - Внутренние узлы – условия на атрибуты
  - Каждая исходящая ветка соответствует выходному значению условия, ветка целиком – альтернативное решение
  - В листьях метки классов (или распределение меток классов) или значения целевой переменной для регрессии
  - Каждому узлу соответствует область в пространстве признаков  $R$
  - Области для листьев – финальные, не содержат внутри других областей
- Построение дерева обычно – 2 фазы
  - Построение: в начале в корне все примеры, далее рекурсивное разбиение множества примеров по выбранному атрибуту
  - «отсечение» ветвей pruning - выявление и удаление ветвей (решений), приводящих к шуму или к выбросам
- Применение дерева решений для нового объекта
  - Проверка атрибутов – путь по ветви до листа. В листе отклик.

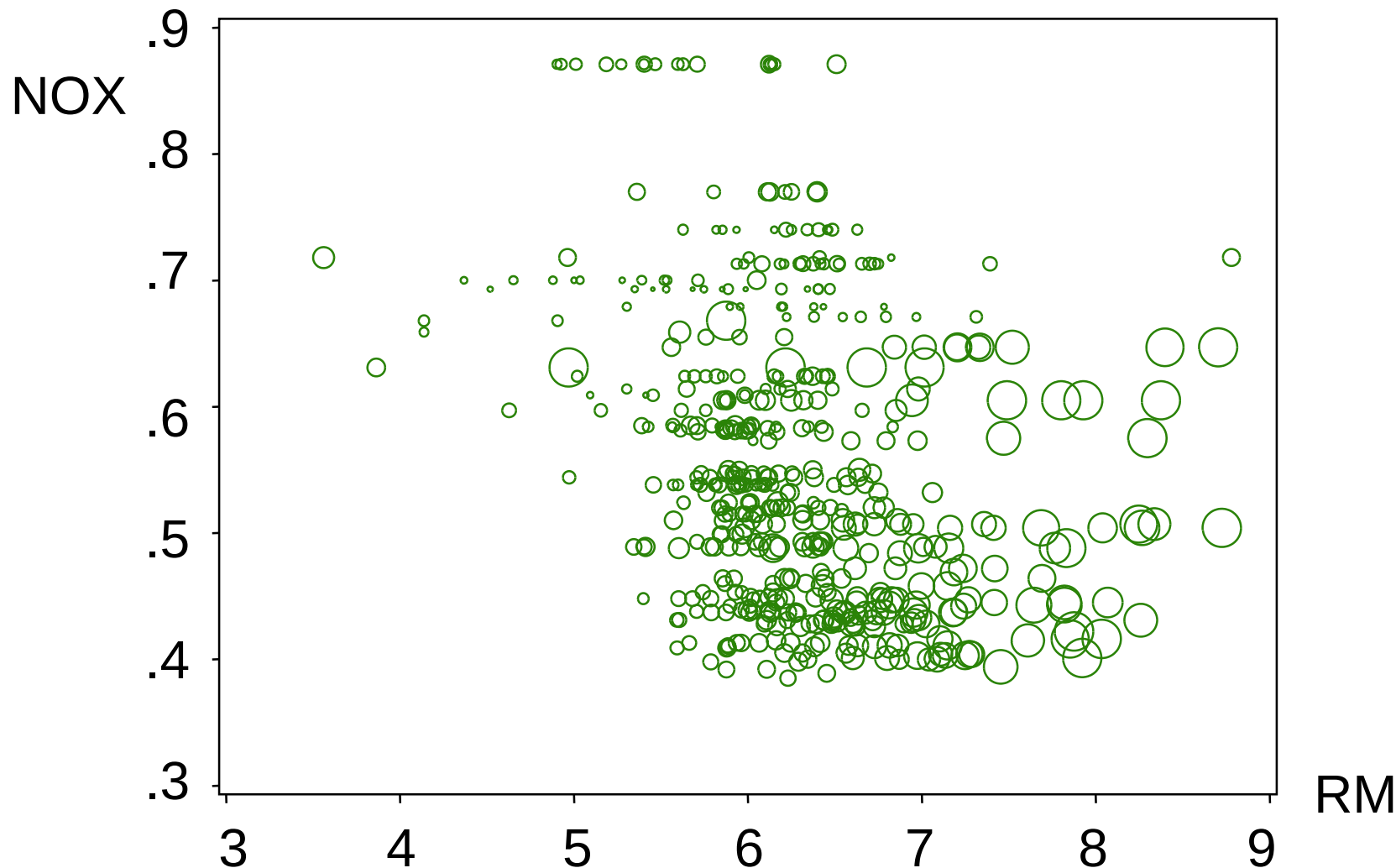
# Пример: категориальный отклик



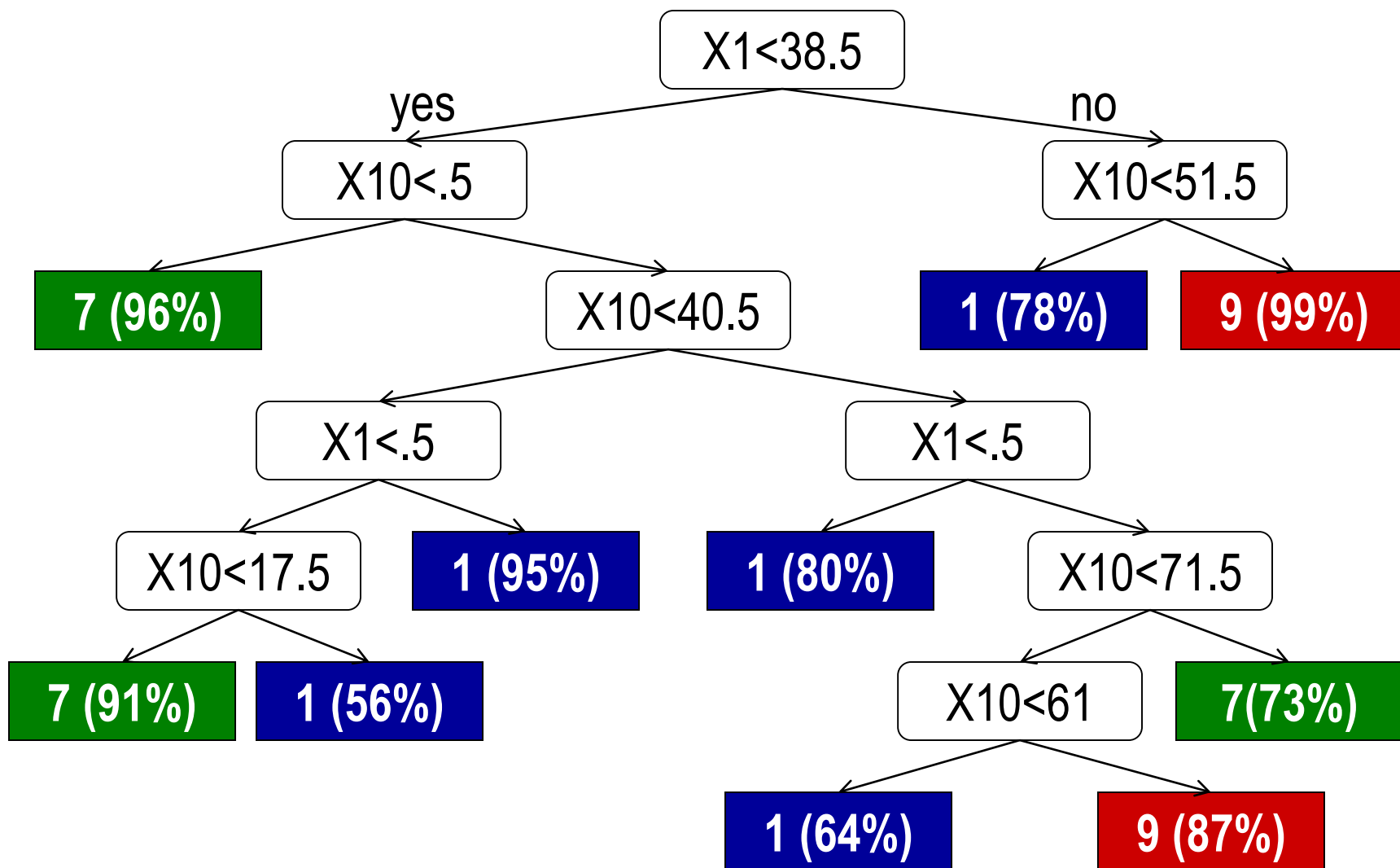
# Пример: непрерывный отклик



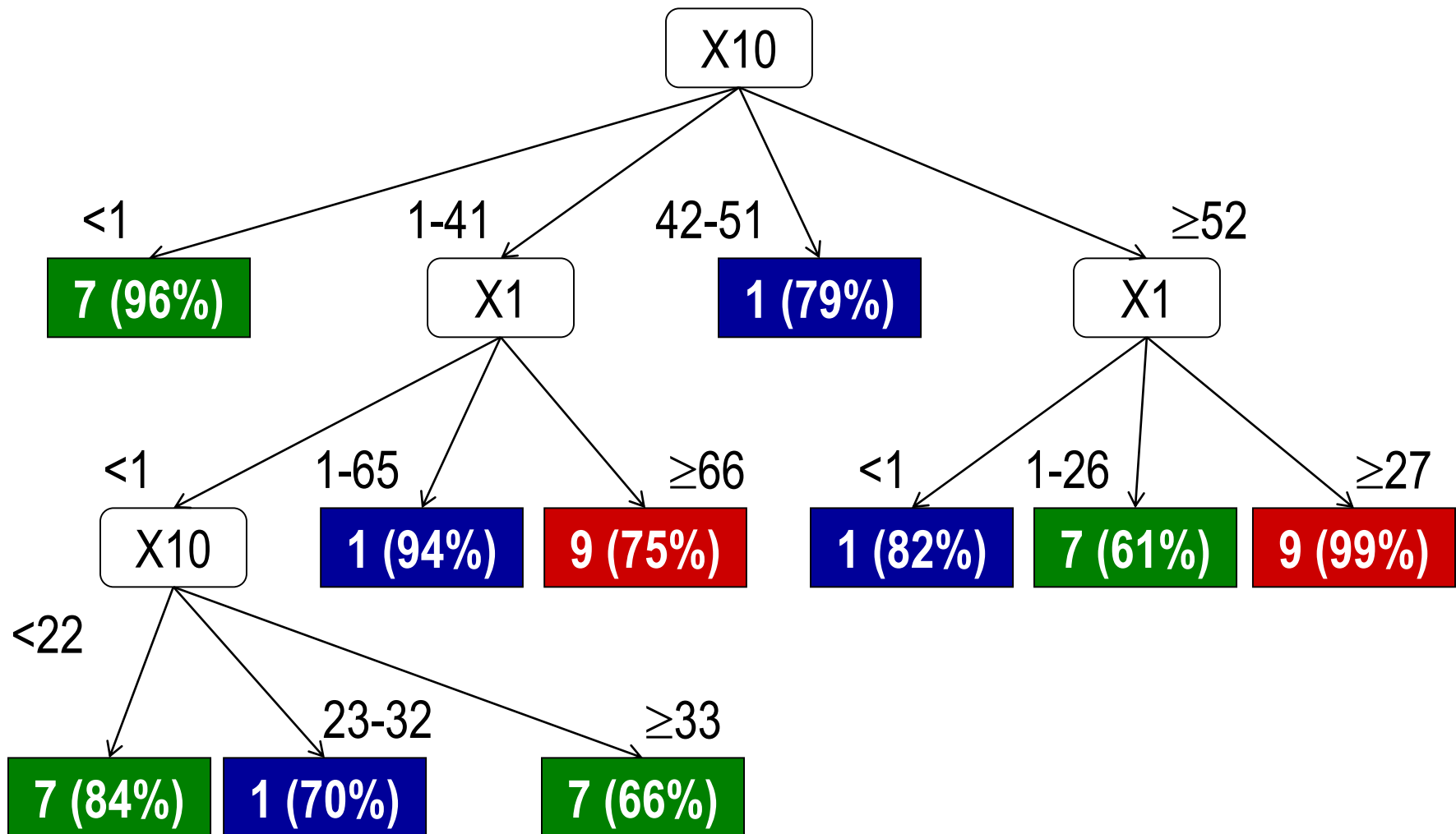
# Проекция непрерывного отклика на пространство признаков



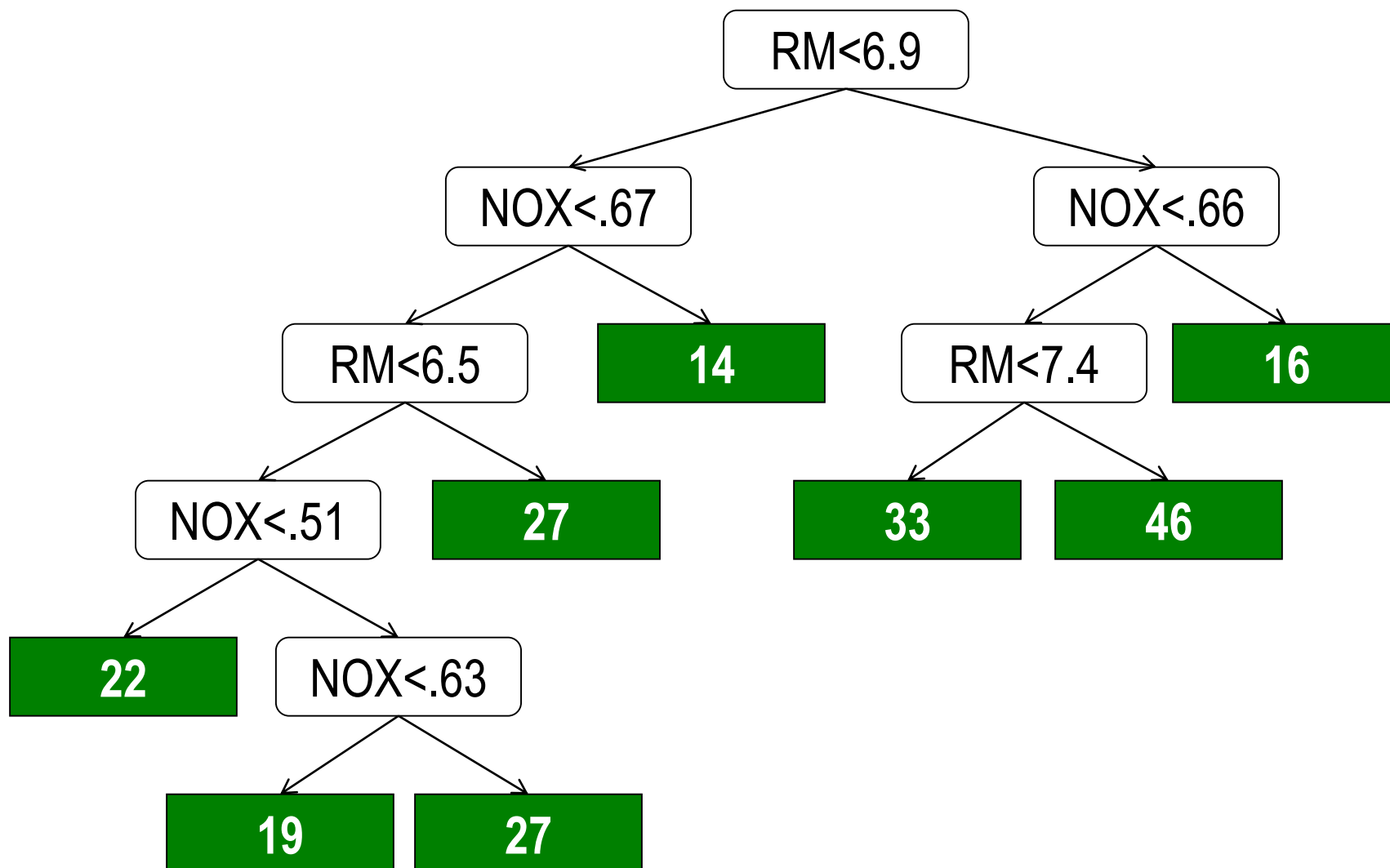
# Дерево решений для классификации



# Множественные разбиения (не бинарное дерево)



# Дерево решений для регрессии

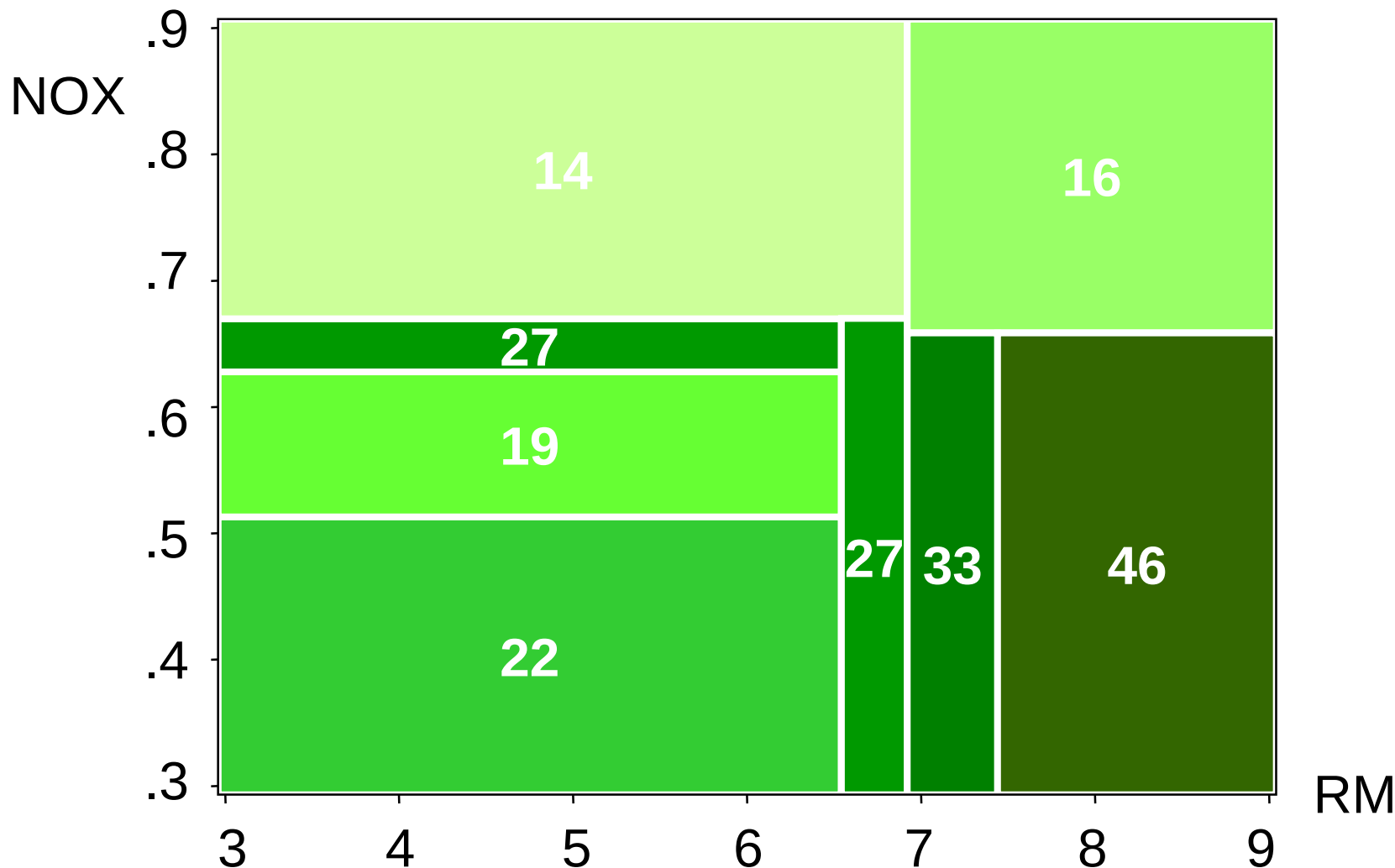


# Листья = Логические правила

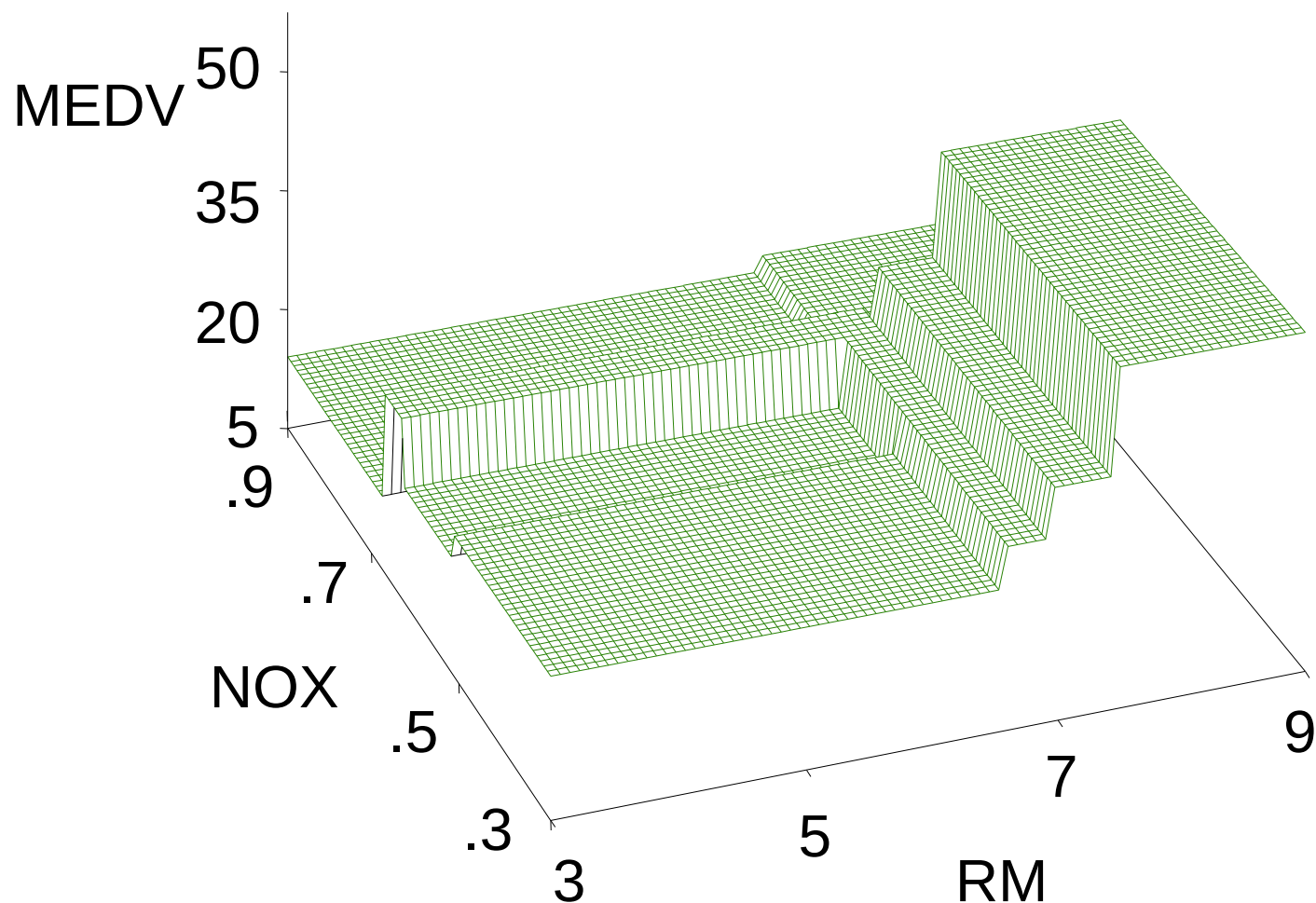
If  $RM \in \{values\}$  and  $NOX \in \{values\}$ , then  $MEDV=value$ .

| <u>Leaf</u> | <u>RM</u>  | <u>NOX</u> | <u>Predicted MEDV</u> |
|-------------|------------|------------|-----------------------|
| 1           | <6.5       | <.51       | 22                    |
| 2           | <6.5       | [.51, .63) | 19                    |
| 3           | <6.5       | [.63, .67) | 27                    |
| 4           | [6.5, 6.9) | <.67       | 27                    |
| 5           | <6.9       | $\geq$ .67 | 14                    |
| 6           | [6.9, 7.4) | <.66       | 33                    |
| 7           | $\geq$ 7.4 | <.66       | 46                    |
| 8           | $\geq$ 6.9 | $\geq$ .66 | 16                    |

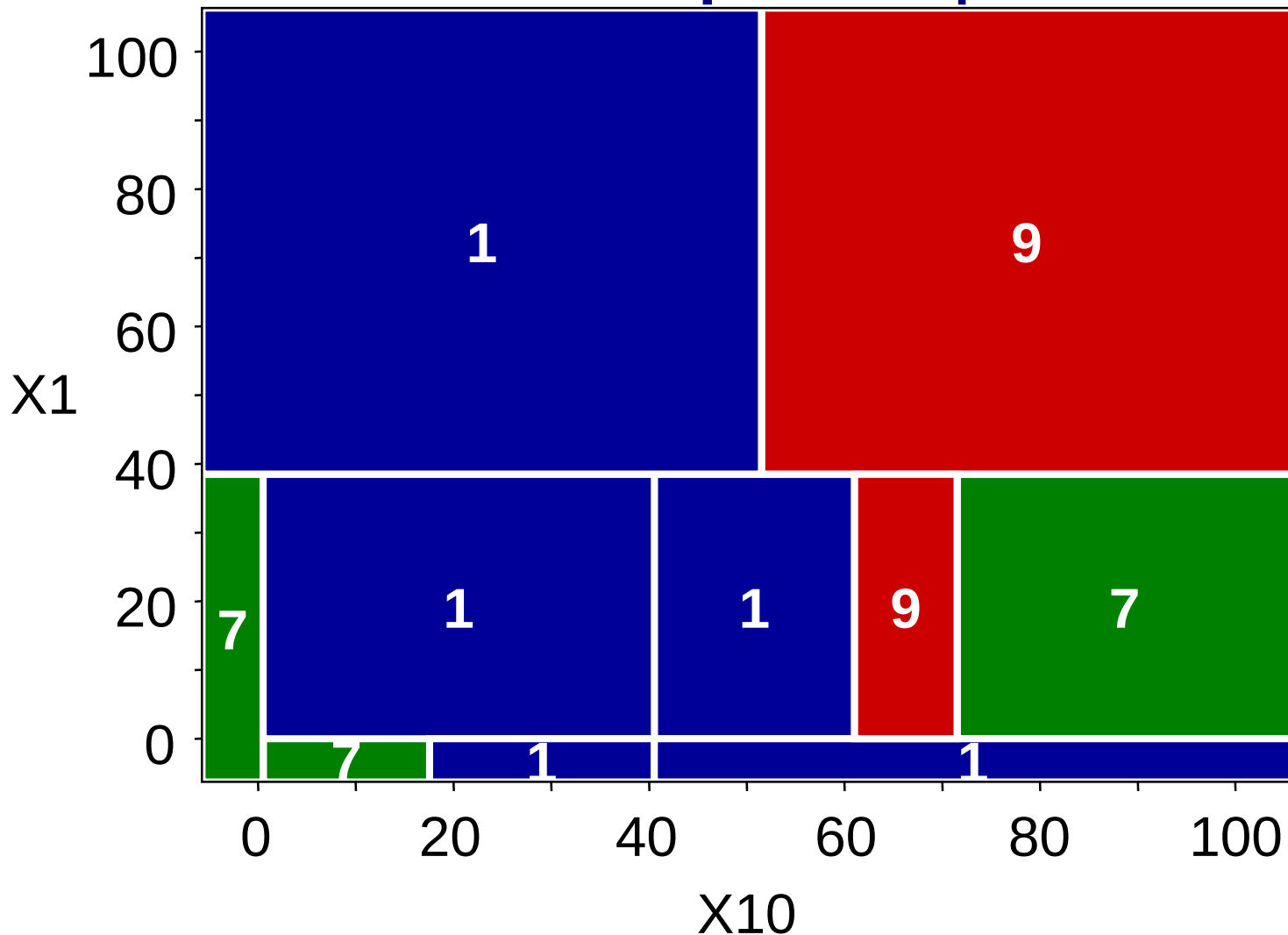
# Разбиение пространства признаков на области



# Многомерная ступенчатая функция



# Регионы решений для классификации



# Листья дерева решений для классификации

| <u>Leaf</u> | <u>Pr(<b>1</b> x)</u> | <u>Pr(<b>7</b> x)</u> | <u>Pr(<b>9</b> x)</u> | <u>Decision</u> |
|-------------|-----------------------|-----------------------|-----------------------|-----------------|
| 1           | .03                   | .96                   | .01                   | <b>7</b>        |
| 2           | .09                   | .91                   | .00                   | <b>7</b>        |
| 3           | .56                   | .44                   | .00                   | <b>1</b>        |
| 4           | .95                   | .05                   | .00                   | <b>1</b>        |
| 5           | .80                   | .10                   | .10                   | <b>1</b>        |
| 6           | .64                   | .09                   | .27                   | <b>1</b>        |
| 7           | .00                   | .13                   | .87                   | <b>9</b>        |
| 8           | .10                   | .73                   | .17                   | <b>7</b>        |
| 9           | .78                   | .01                   | .21                   | <b>1</b>        |
| 10          | .01                   | .00                   | .99                   | <b>9</b>        |

# Процесс построения деревьев решений

- Разделяем пространство признаков (то есть, набор возможных значений для  $X_1, X_2, \dots, X_p$ ) в  $J$  отдельных и непересекающихся областей  $R_1, R_2, \dots, R_J$ . Таким образом, чтобы поведение отклика в каждой области было «однородным».
- Теоретически, области могут иметь любую форму. Тем не менее, чаще всего используются многомерные *прямоугольники* для простоты и удобства интерпретации полученной модели.
- Цель состоит в том, чтобы найти прямоугольные области  $R_1, \dots, R_J$ , так чтобы минимизировать некий целевой критерий – *критерий разбиения*.

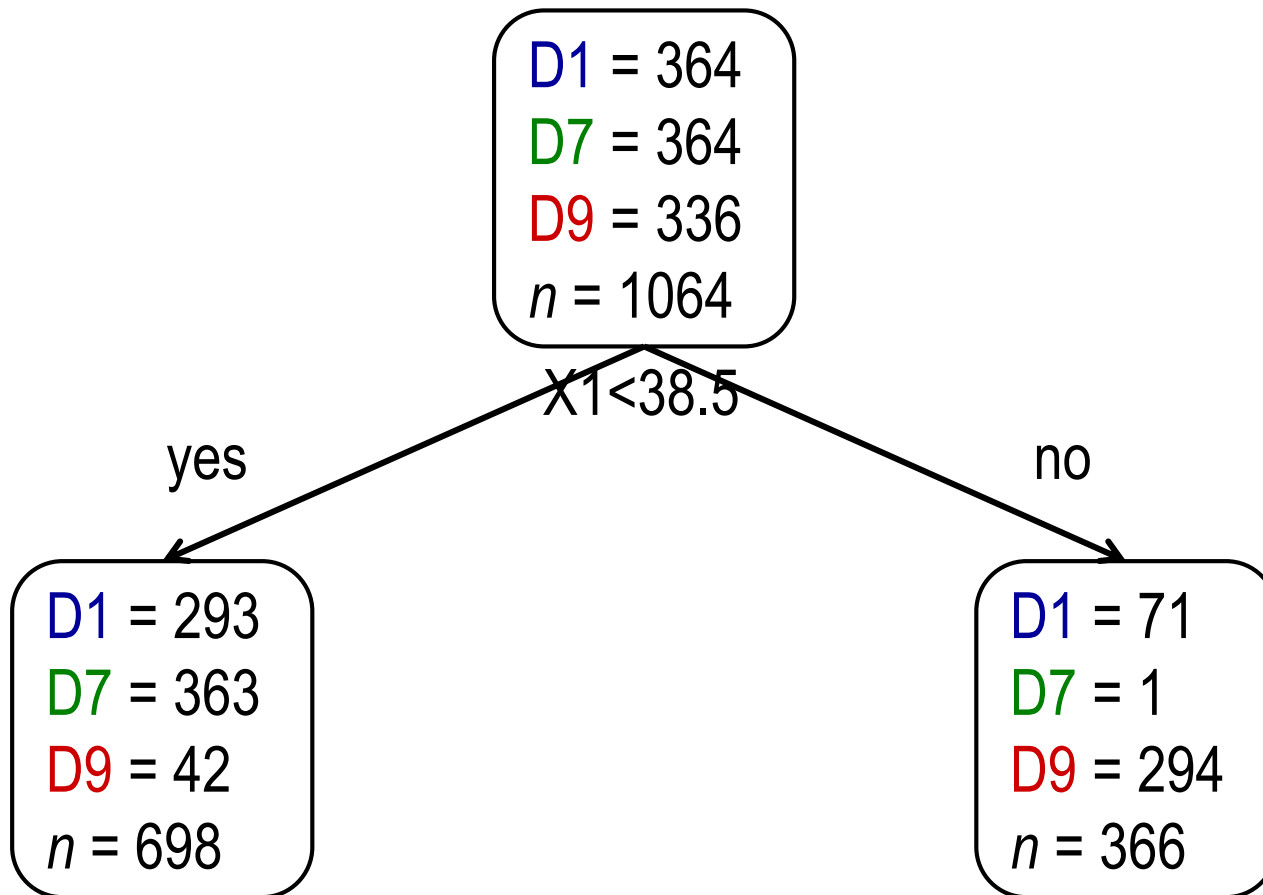
# Процесс построения деревьев решений

- Вычислительно нецелесообразно рассматривать все возможные разбиения пространства признаков в  $J$  областей.
- По этой причине используется *нисходящий, жадный* подход, известный как рекурсивное разделение.
- Подход называется *нисходящий*, потому что он начинается в верхней части дерева, а затем последовательно разбивает пространство признаков; каждое разбиение приводит к образованию новых ветвей, расположенных ниже по дереву.
- Подход называется *жадным*, потому что на каждом этапе процесса построения деревьев *лучшее* разбиение осуществляется на конкретном шаге, вместо того, чтобы просматривать дальше и выбирать разбиение, которое приведет к лучшему дереву в дальнейшем.

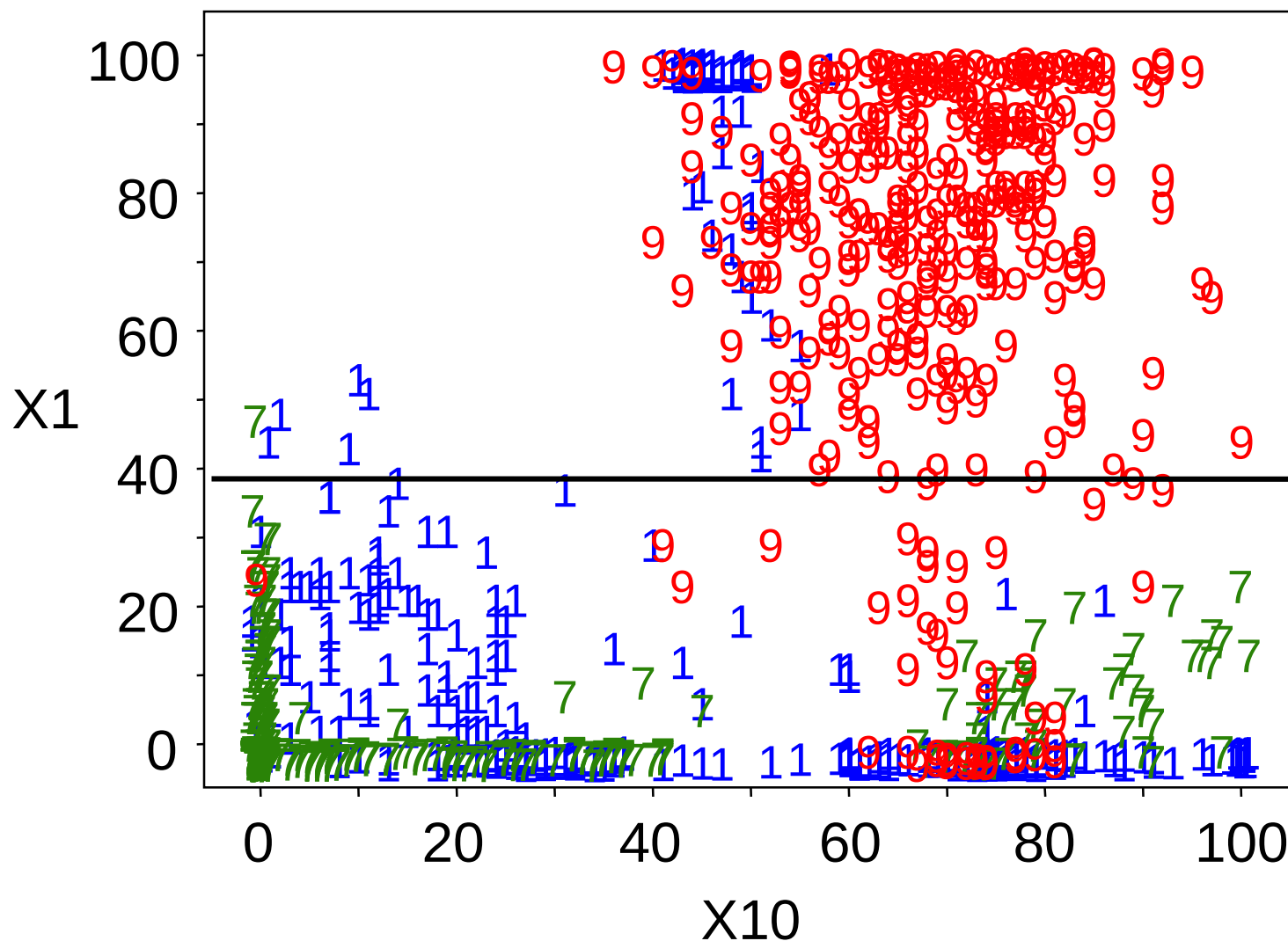
# Процесс построения деревьев решений

- Сначала выбирается переменная  $X_j$  и точка разбиения  $s$ , так что разбиение пространства признаков на области  $\{X|X_j < s\}$  и  $\{X|X_j \geq s\}$  приводит к максимально возможному улучшению критерия (для бинарного дерева).
- Затем повторяется этот процесс и ищется лучшая переменная и точка разбиения в каждой из результирующих областей.
- Однако на этот раз, вместо разделения всего пространства признаков, мы разделяем одну из ранее выделенных областей.
- Процесс продолжается до тех пор, пока не будет достигнут критерий останова; например, мы можем продолжать процесс до тех пор, пока не останется областей, содержащих более пяти наблюдений.

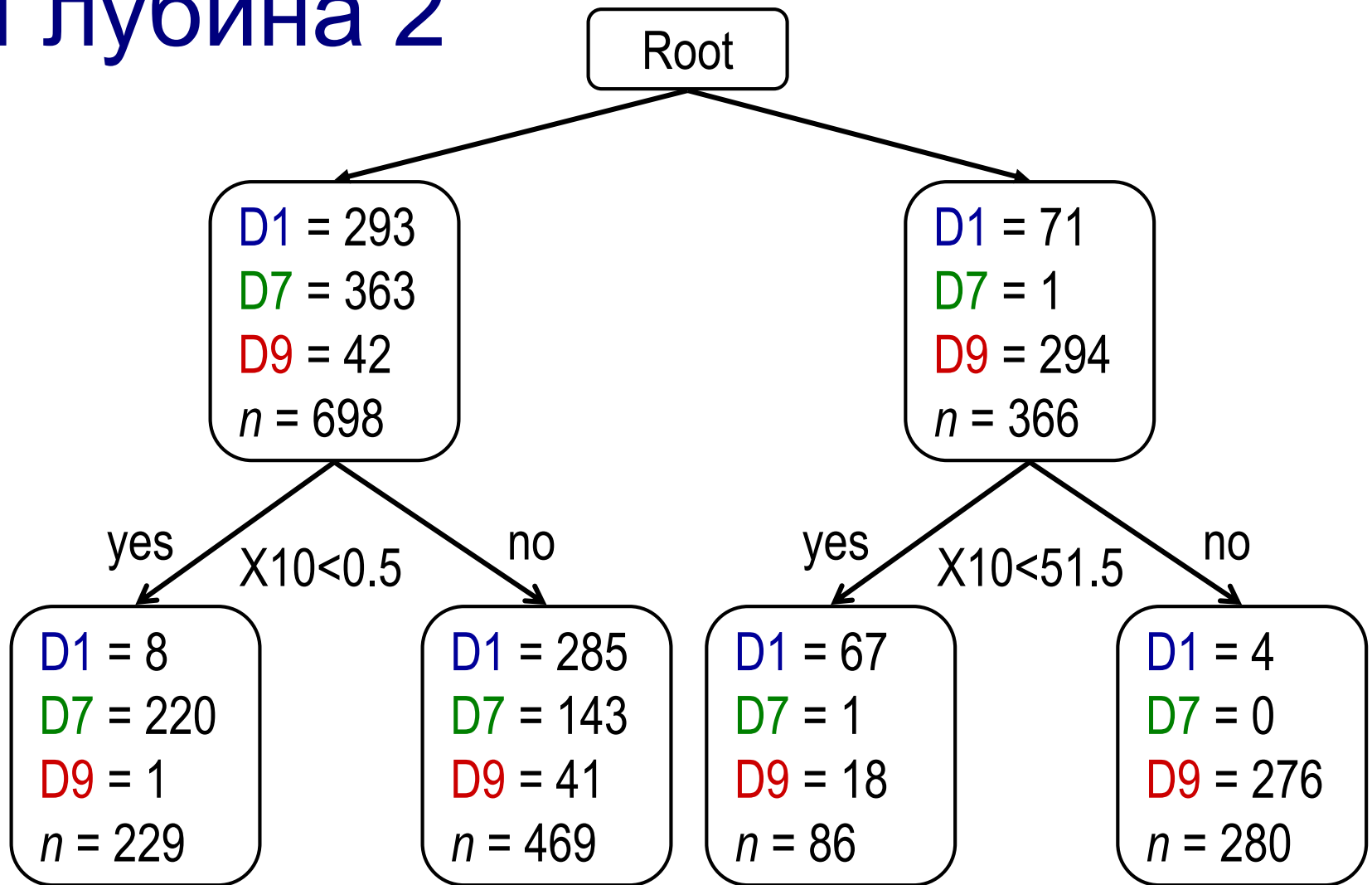
# Шаг разбиения узла



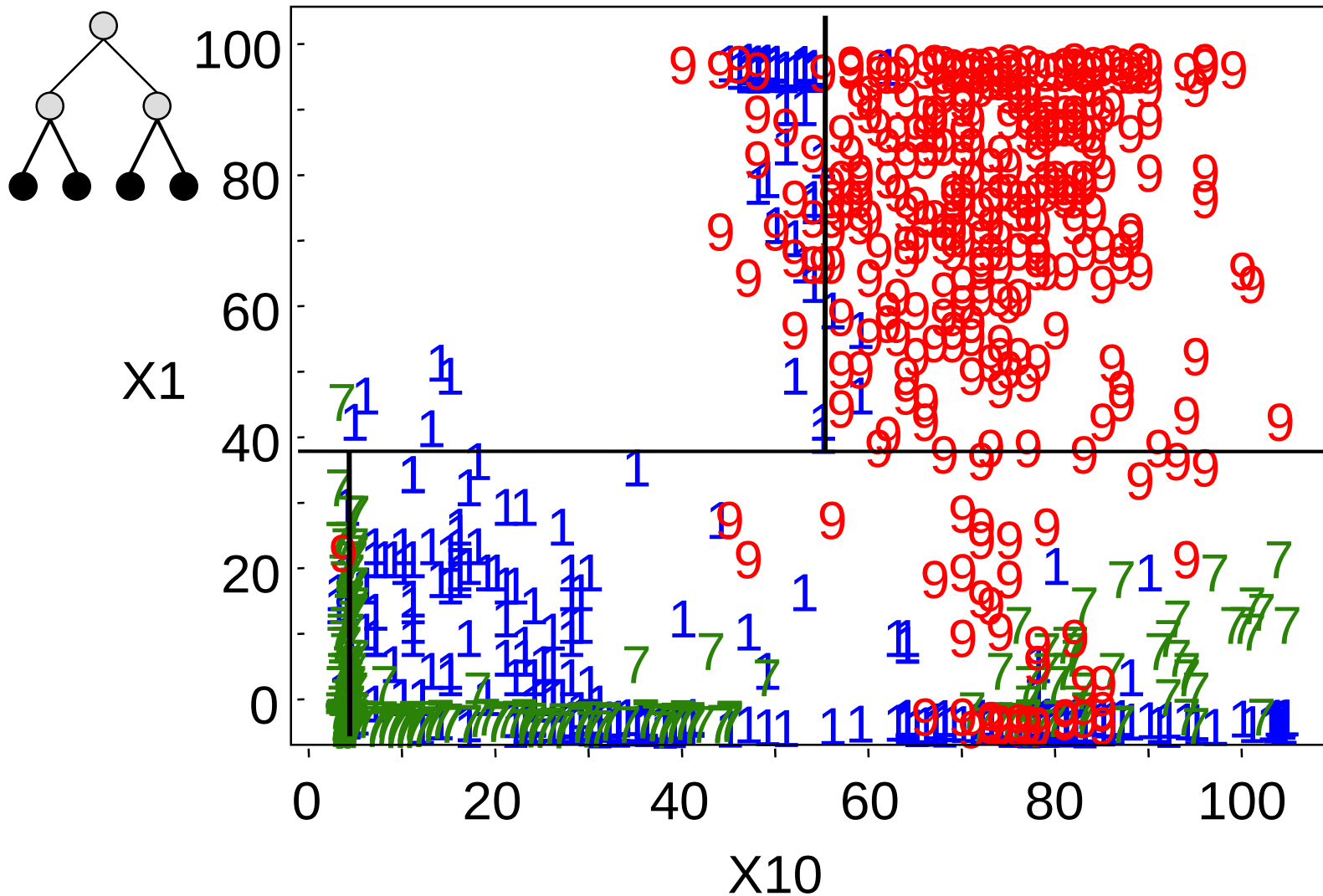
# Разбиение деревом глубины 1



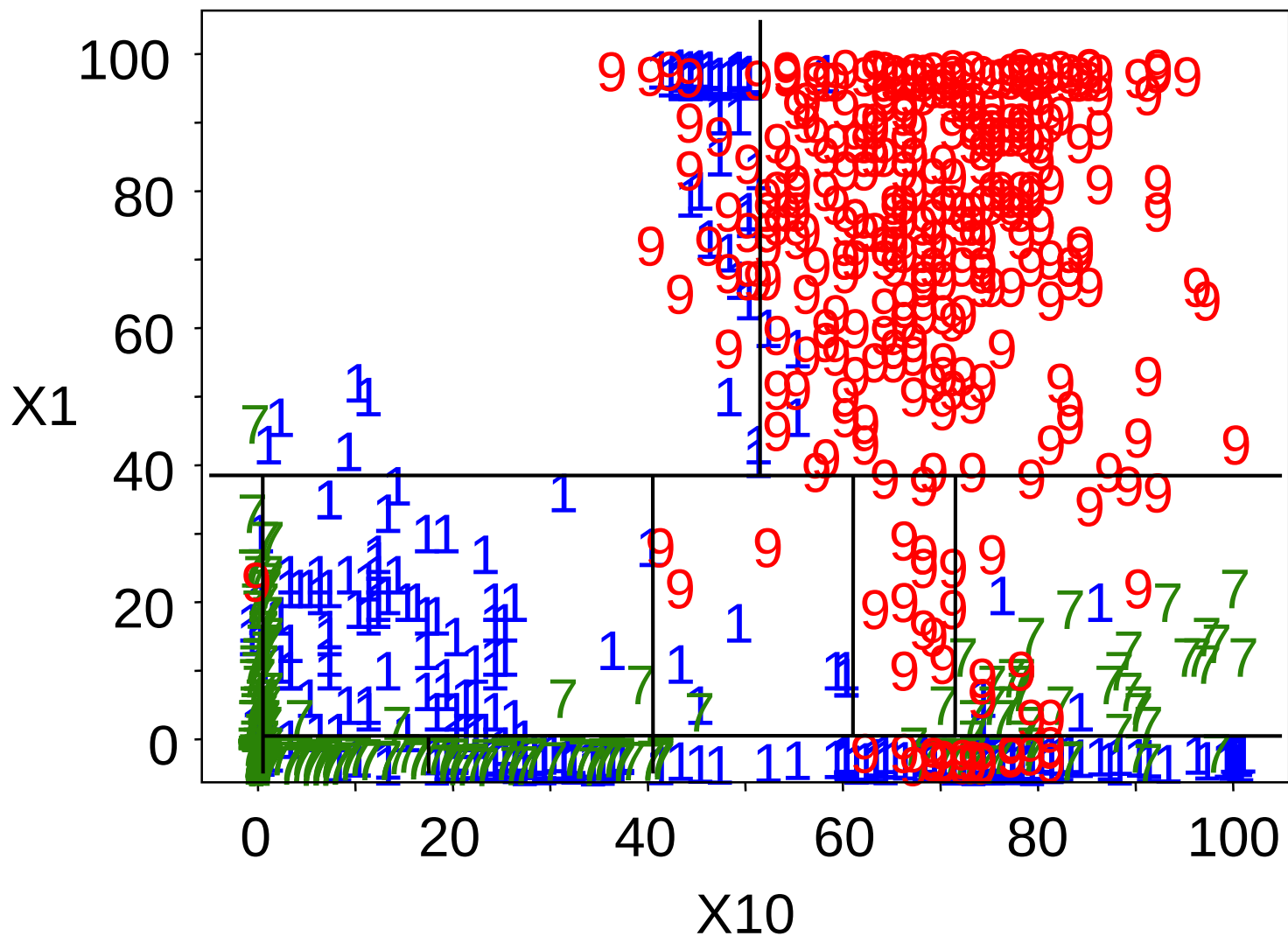
# Глубина 2



# Разбиение деревом глубины 2



# Финальное разбиение



# Рассмотрим

- Поиск разбиения по переменным

- ☐ Ординальным
- ☐ Категориальным

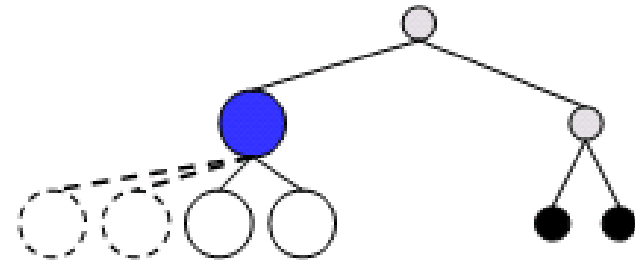
- Множественные разбиения

- Критерии разбиения

- ☐ Уменьшение разнородности
- ☐ Хи2 тест

- Регрессионные деревья

- Пропущенные значения



| <u>Variable</u> | <u>Values</u> |
|-----------------|---------------|
| X10             | 0.5           |
| X10             | 1.8           |
| X10             | 11, 46        |
| X1              | 2.4           |
| X1              | 1, 4, 61      |
| :               | :             |
| :               | :             |
| :               | :             |

# Разбиение по ординальной переменной

## Splits

1—234

12—34

123—4

$$\begin{pmatrix} 3 \\ 1 \end{pmatrix} = 3$$

$$\begin{pmatrix} L-1 \\ B-1 \end{pmatrix} = \frac{(L-1)!}{(B-1)!(L-B)!}$$

1—2—34

1—23—4

12—3—4

$$\begin{pmatrix} 3 \\ 2 \end{pmatrix} = 3$$

$$\sum_{l=2}^L \begin{pmatrix} L-1 \\ l-1 \end{pmatrix} = 2^{L-1} - 1$$

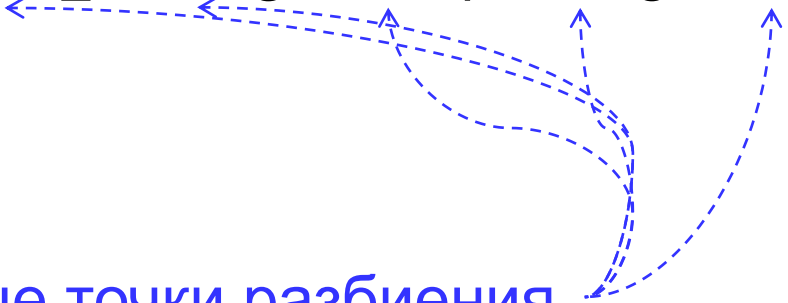
1—2—3—4

$$\begin{pmatrix} 3 \\ 3 \end{pmatrix} = 1$$

# Ординальные переменные

|          |      |     |     |     |     |      |
|----------|------|-----|-----|-----|-----|------|
| X        | .20  | 1.7 | 3.3 | 3.5 | 14  | 2515 |
| $\ln(X)$ | -1.6 | .53 | 1.2 | 1.3 | 2.6 | 7.8  |
| rank(X)  | 1    | 2   | 3   | 4   | 5   | 6    |

Потенциальные точки разбиения



# Разбиение категориальной переменной

$$S(L, B) = B \cdot S(L-1, B) + S(L-1, B-1)$$

1—234  
 2—134  
 3—124  
 4—123  
 12—34  
 13—24  
 14—23  
 1—2—34  
 1—3—24  
 1—4—23  
 2—3—14  
 2—4—13  
 3—4—12  
 1—2—3—4

| <i>L</i> | <i>B</i> : | 2   | 3    | 4    | total |
|----------|------------|-----|------|------|-------|
|          |            |     |      |      |       |
|          | 2          | 1   |      |      | 1     |
|          | 3          | 3   | 1    |      | 4     |
|          | 4          | 7   | 6    | 1    | 14    |
|          | 5          | 15  | 25   | 10   | 51    |
|          | 6          | 31  | 90   | 65   | 202   |
|          | 7          | 63  | 301  | 350  | 876   |
|          | 8          | 127 | 966  | 1701 | 4139  |
|          | 9          | 255 | 3025 | 7770 | 21146 |

# Основные моменты поиска разбиения

- Только бинарное разбиение

- ординальные =  $L - 1$
  - категориальные =  $2^{L-1} - 1$

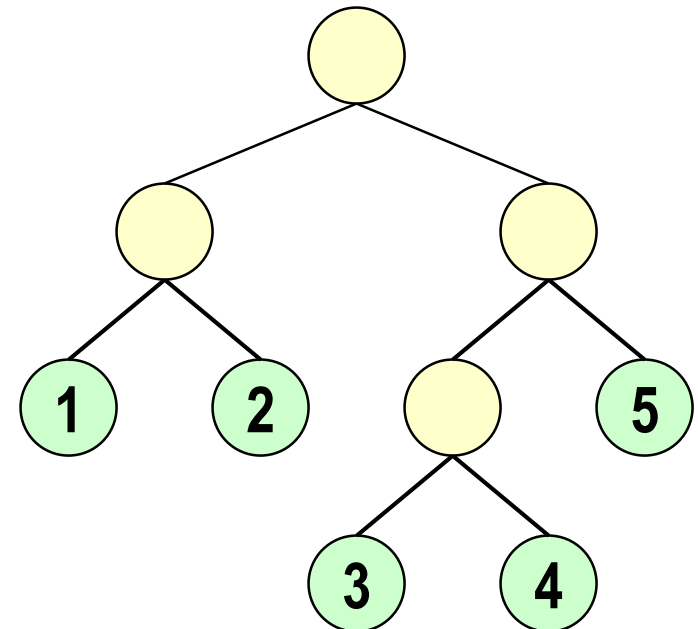
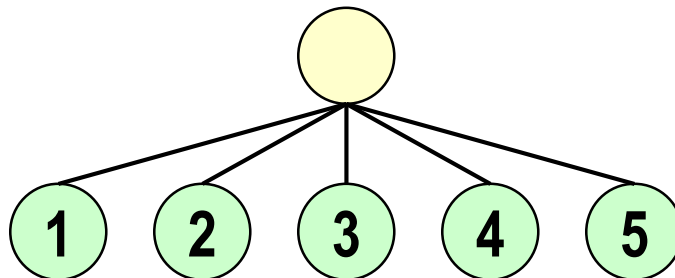
- Агломеративная кластеризация

- Вариантов разбиения

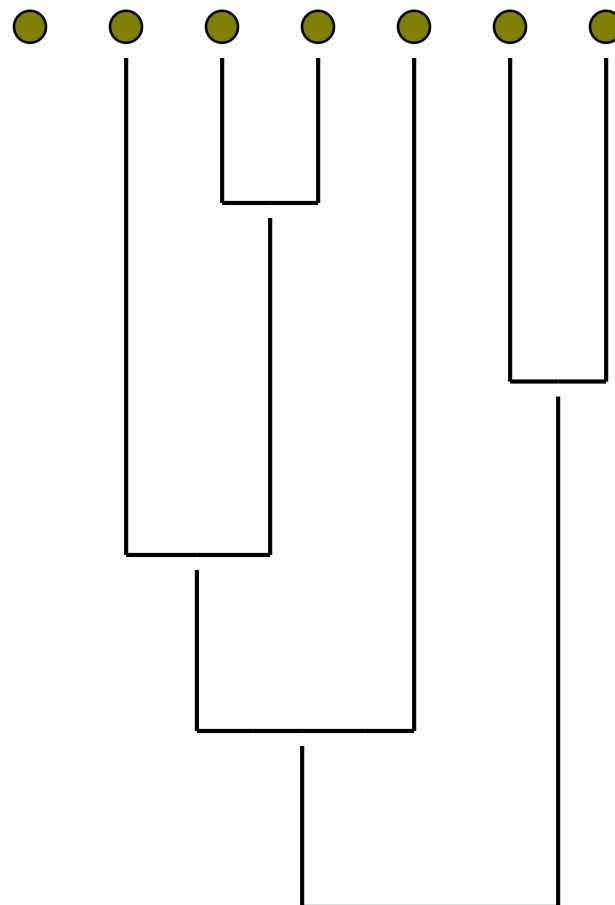
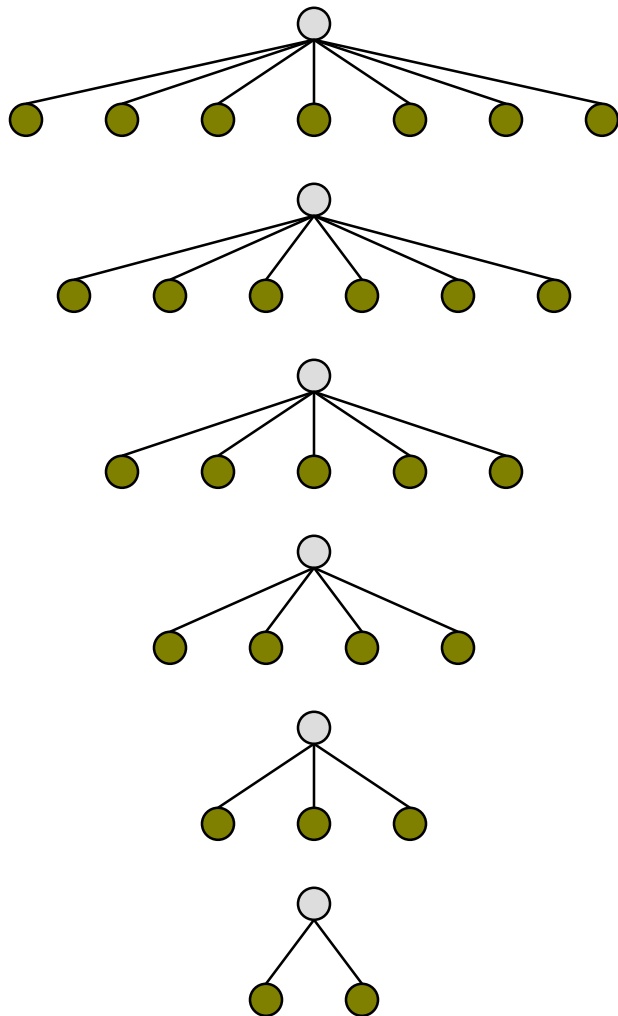
- Kass (1980)

- Минимальный размер листа

- Множественное или бинарное



# Кластеризация гипотез



# Критерии разбиения для категориального отклика

X1: <38.5    ≥38.5

|          |            |            |
|----------|------------|------------|
| <b>1</b> | <b>293</b> | <b>71</b>  |
| <b>7</b> | <b>363</b> | <b>1</b>   |
| <b>9</b> | <b>42</b>  | <b>294</b> |

ΔGini    Δentropy    logworth

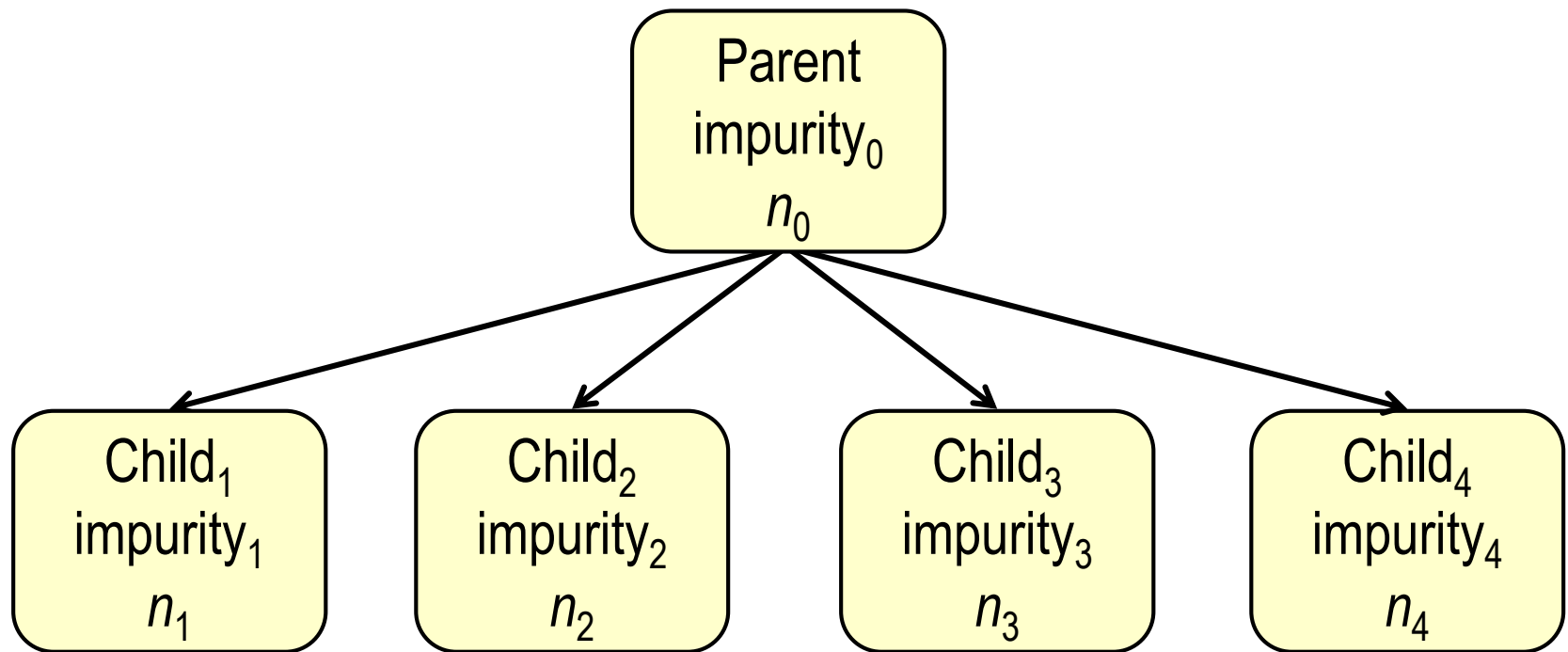
.197       .504       140

X10: <0.5    1-41    42-51    ≥51.5

|          |            |            |           |            |
|----------|------------|------------|-----------|------------|
| <b>1</b> | <b>9</b>   | <b>143</b> | <b>65</b> | <b>147</b> |
| <b>7</b> | <b>221</b> | <b>88</b>  | <b>1</b>  | <b>54</b>  |
| <b>9</b> | <b>1</b>   | <b>4</b>   | <b>16</b> | <b>315</b> |

.255       .600       172

# Уменьшение разнородности

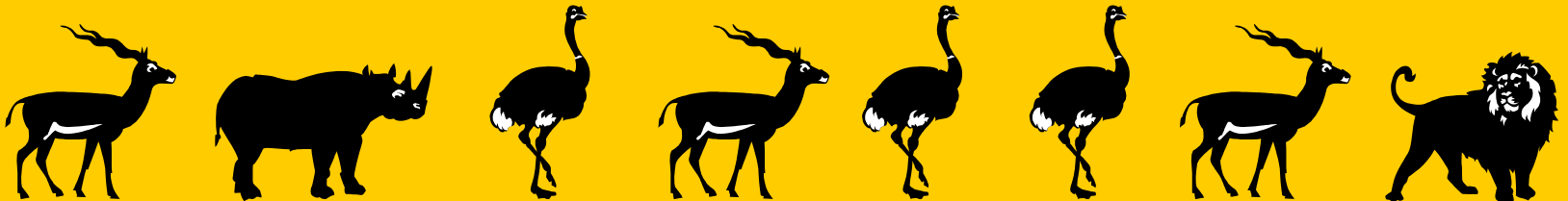


$$\Delta i = i(0) - \left( \frac{n_1}{n_0} i(1) + \frac{n_2}{n_0} i(2) + \frac{n_3}{n_0} i(3) + \frac{n_4}{n_0} i(4) \right)$$

# Индекс Джини

$$1 - \sum_{j=1}^r p_j^2 = 2 \sum_{j < k} p_j p_k$$

**Высокая вариация, низкая чистота**



$$\text{Pr(interspecific encounter)} = 1 - 2(3/8)^2 - 2(1/8)^2 = .69$$

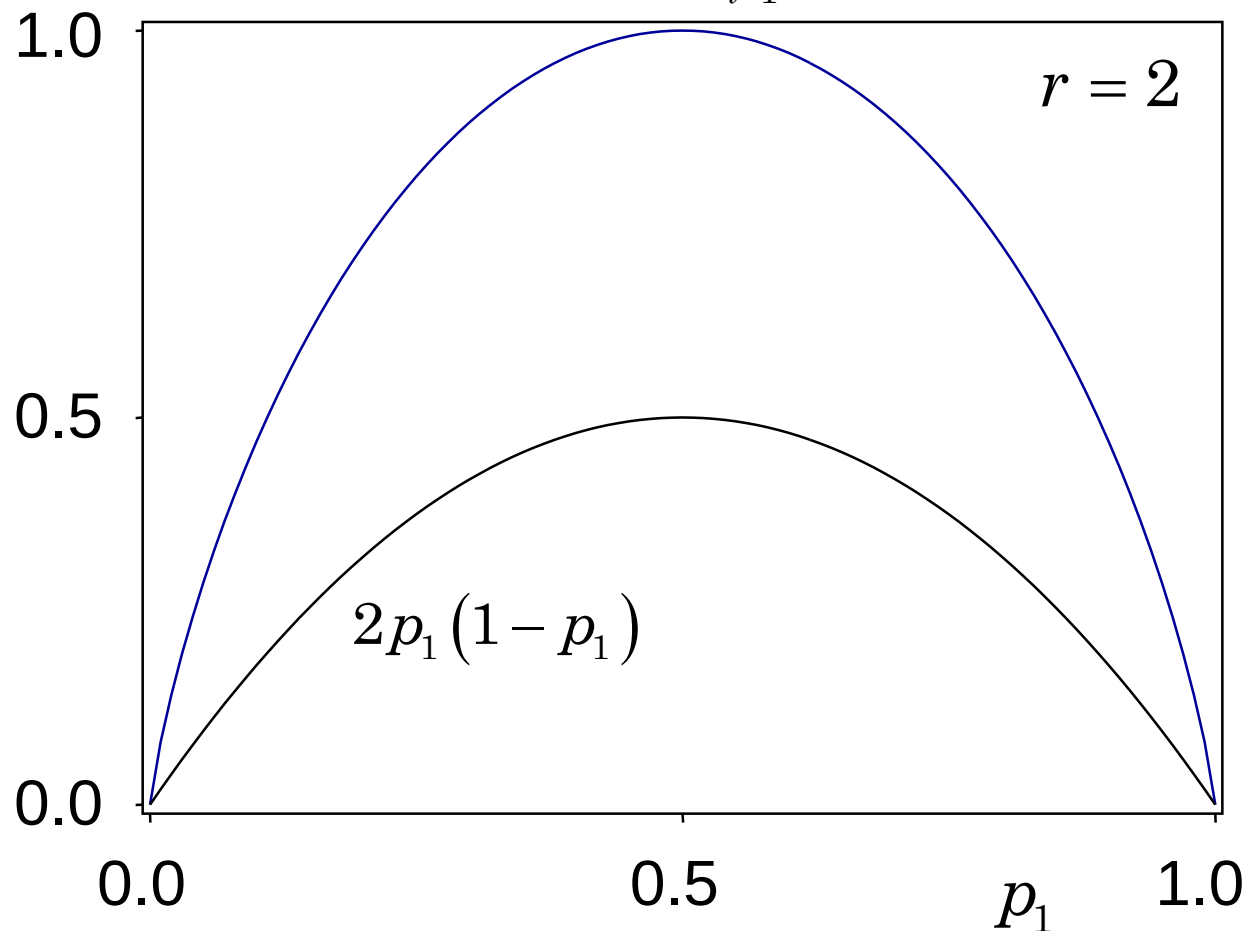
**Низкая вариация, высокая чистота**



$$\text{Pr(interspecific encounter)} = 1 - (6/7)^2 - (1/7)^2 = .24$$

# Энтропия

$$H(p_1, p_2, \dots, p_r) = -\sum_{i=1}^r p_i \log_2(p_i)$$



# Хи2 тест

Наблюдаемые  
частоты

Ожидаемые  
(по гипотезе)

$$\frac{(O - E)^2}{E}$$

X1: <38.5 ≥ 38.5

|          |            |            |          |
|----------|------------|------------|----------|
| <b>1</b> | <b>293</b> | <b>71</b>  | .342     |
| <b>7</b> | <b>363</b> | <b>1</b>   | .342     |
| <b>9</b> | <b>42</b>  | <b>294</b> | .316     |
|          | .656       | .344       | $n=1064$ |

|            |            |
|------------|------------|
| <b>239</b> | <b>125</b> |
| <b>239</b> | <b>125</b> |
| <b>225</b> | <b>116</b> |

|            |            |
|------------|------------|
| <b>12</b>  | <b>23</b>  |
| <b>64</b>  | <b>123</b> |
| <b>149</b> | <b>273</b> |

$$\chi^2_v = \sum \frac{(O - E)^2}{E} = 644$$

$$v = (3 - 1)(2 - 1) = 2$$

# Корректировка $p$ -Value

X1: 38.5

|   |     |     |
|---|-----|-----|
| 1 | 293 | 71  |
| 7 | 363 | 1   |
| 9 | 42  | 294 |

| $\chi^2_\nu$ | $\nu$ | $-\log_{10}(P)$ | $m$ | $-\log_{10}(mP)$ |
|--------------|-------|-----------------|-----|------------------|
| 644          | 2     | 140             | 96  | 138              |

X1: 17.5 36.5

|   |     |    |     |
|---|-----|----|-----|
| 1 | 249 | 42 | 73  |
| 7 | 338 | 25 | 1   |
| 9 | 26  | 16 | 294 |

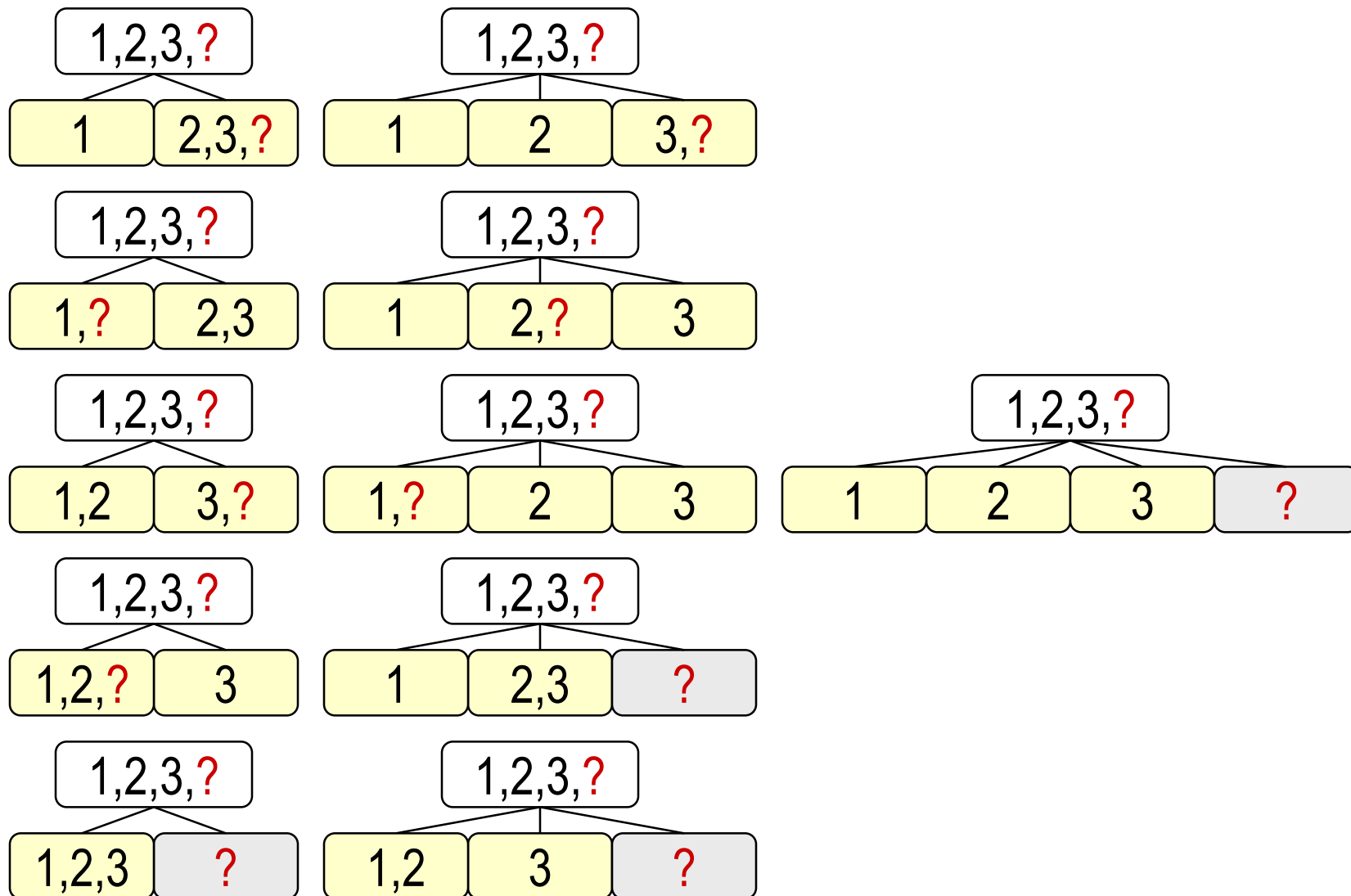
|     |   |     |      |     |
|-----|---|-----|------|-----|
| 660 | 4 | 141 | 4560 | 137 |
|-----|---|-----|------|-----|

X10: 0.5 41.5 51.5

|   |     |     |    |     |
|---|-----|-----|----|-----|
| 1 | 9   | 143 | 65 | 147 |
| 7 | 221 | 88  | 1  | 54  |
| 9 | 1   | 4   | 16 | 315 |

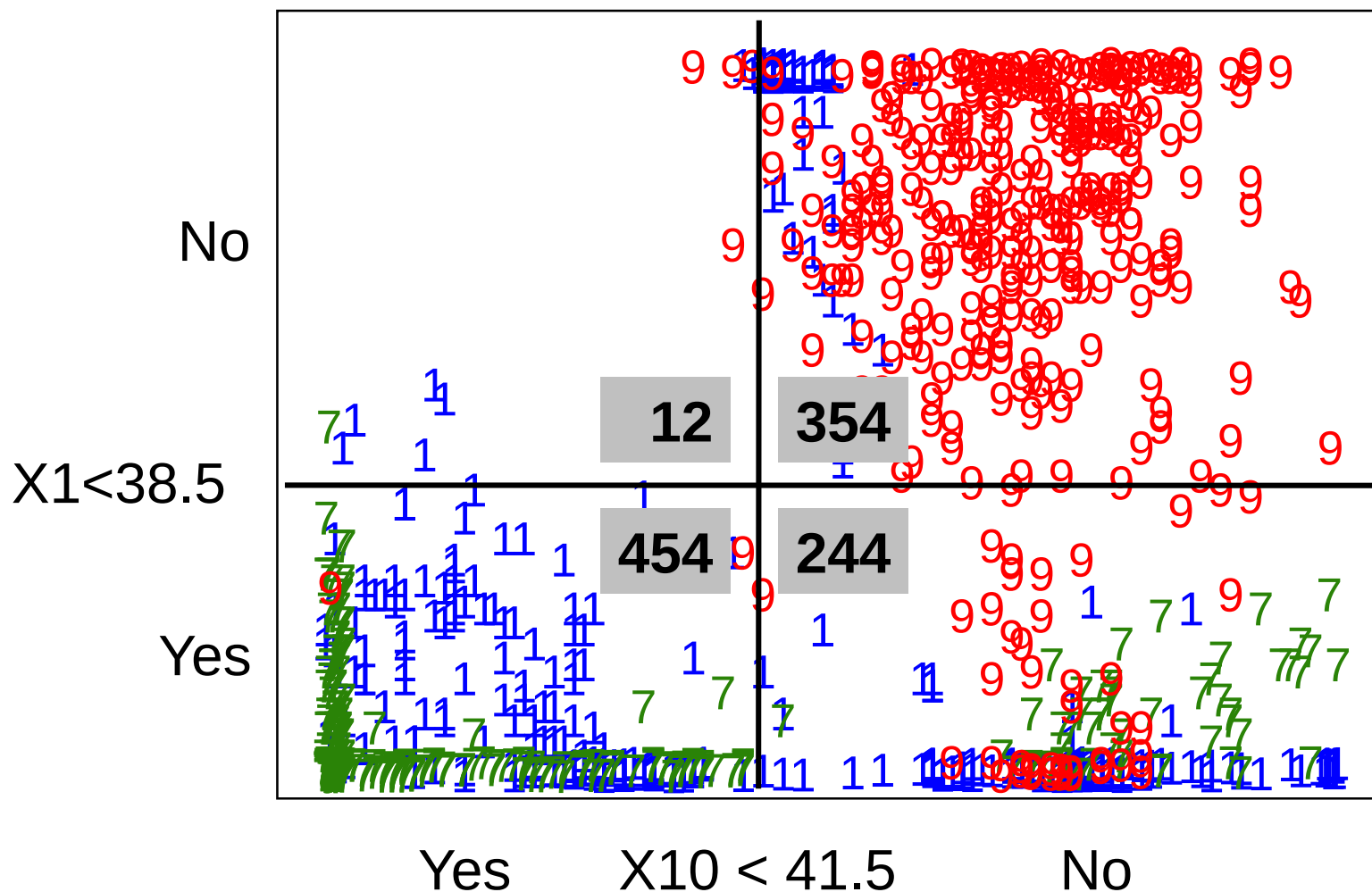
|     |   |     |        |     |
|-----|---|-----|--------|-----|
| 814 | 6 | 172 | 156849 | 167 |
|-----|---|-----|--------|-----|

# Пропущенные значения

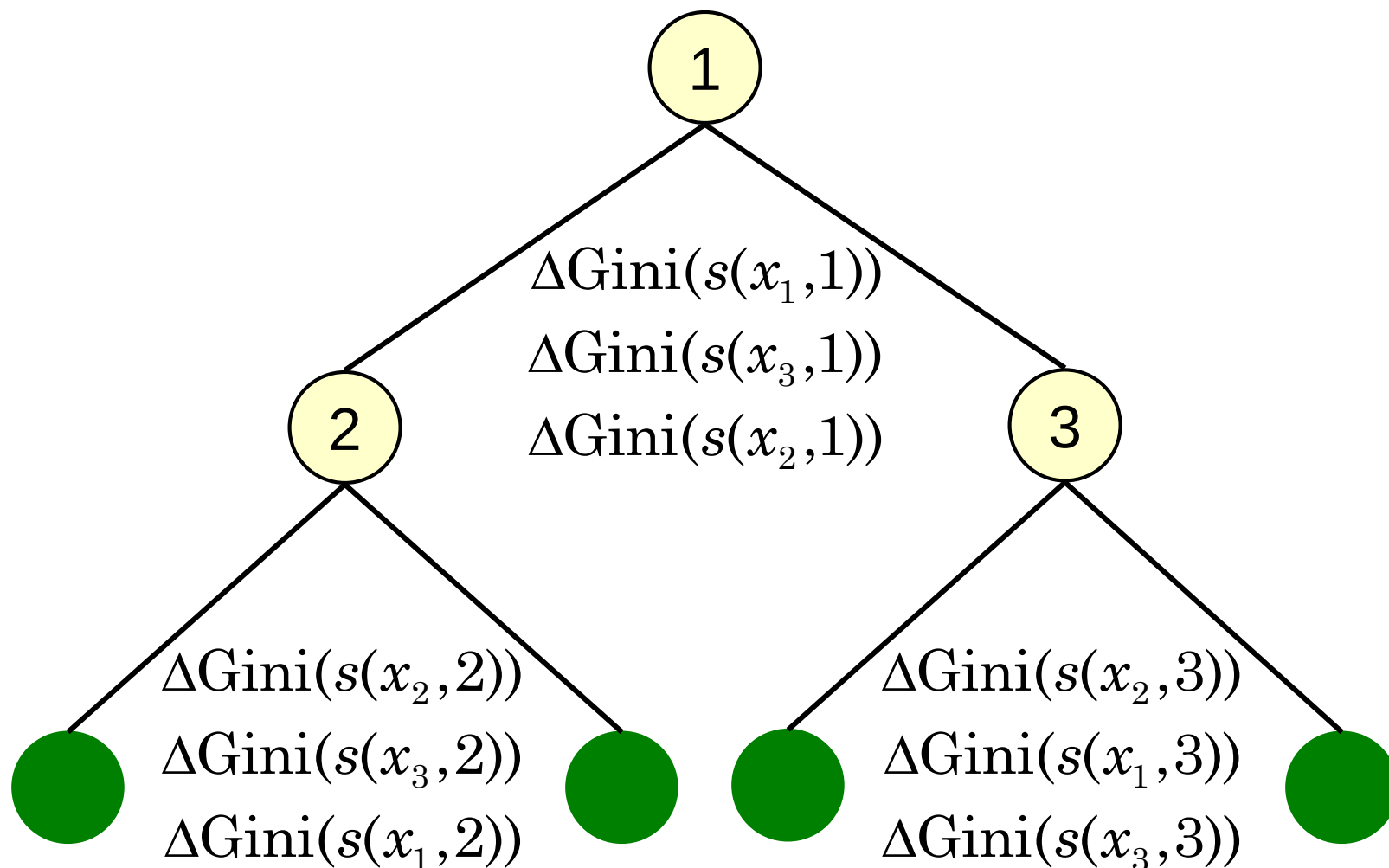


# Суррогатные разбиения

Уровень согласия=76%

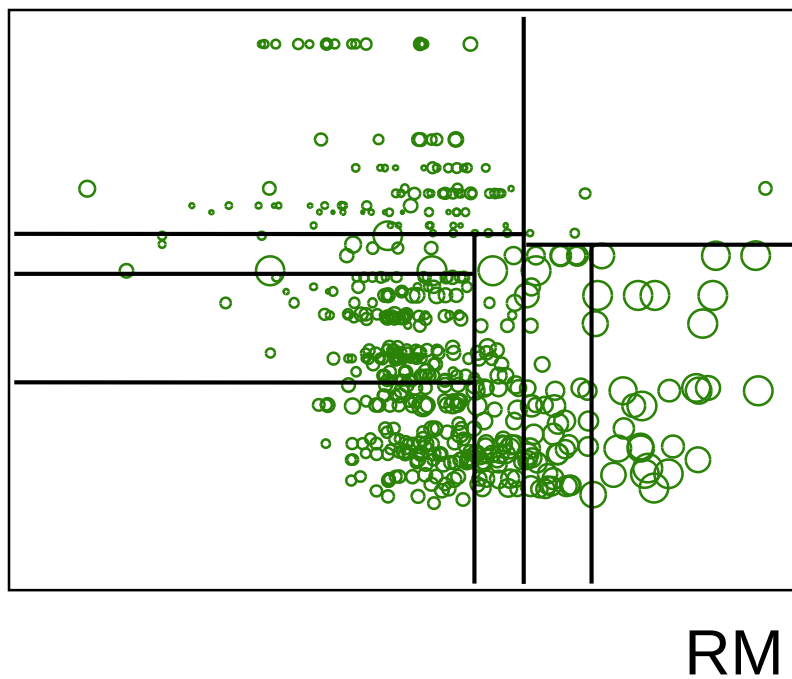


# Важность переменных

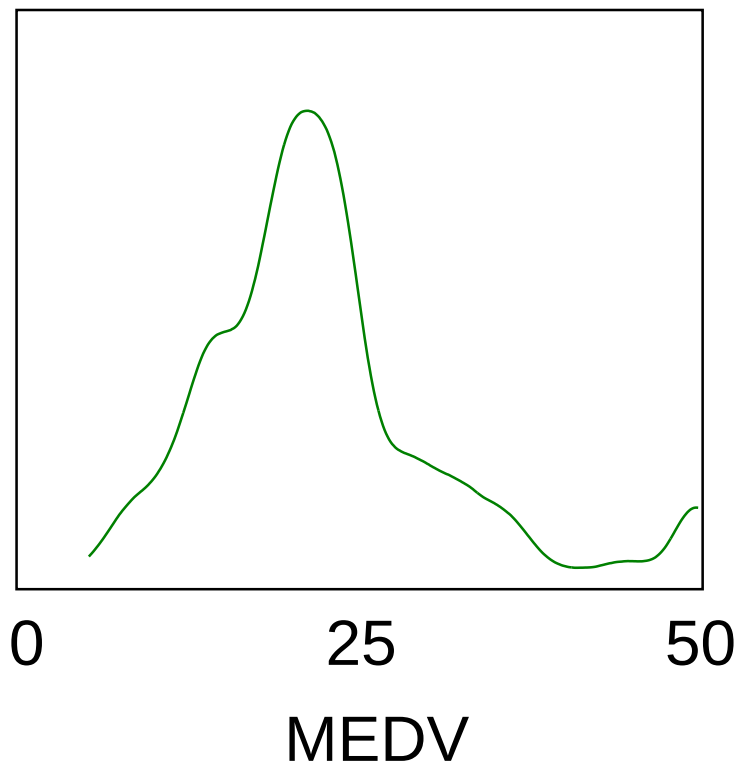


# Непрерывный отклик

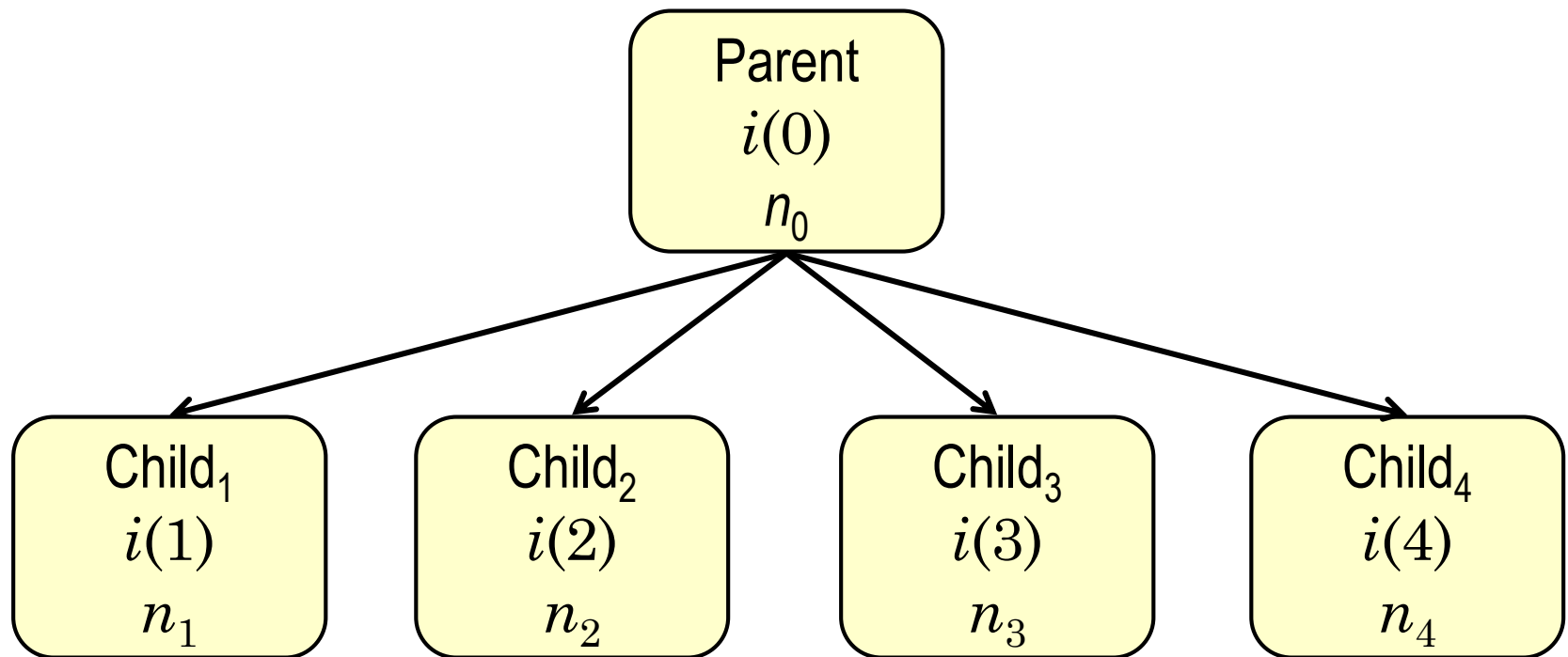
NOX



Плотность

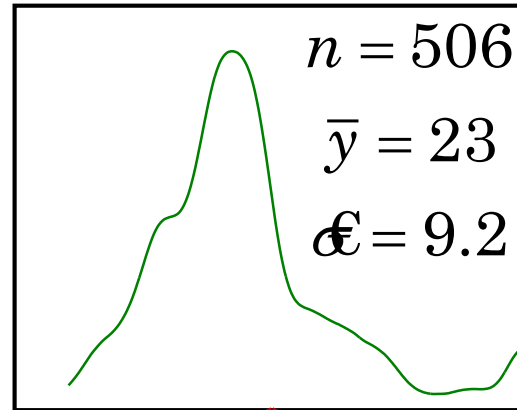
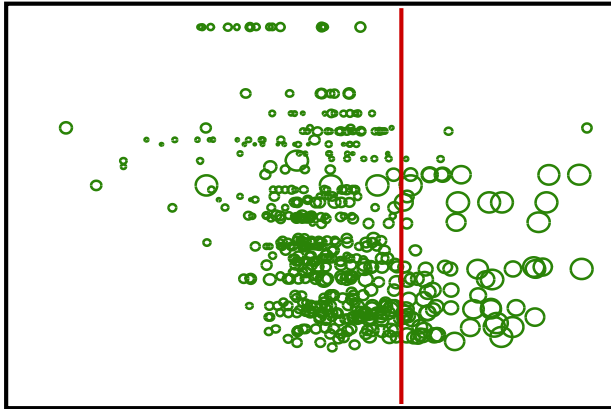


# Уменьшение вариации



$$\Delta i = i(0) - \left( \frac{n_1}{n_0} i(1) + \frac{n_2}{n_0} i(2) + \frac{n_3}{n_0} i(3) + \frac{n_4}{n_0} i(4) \right)$$

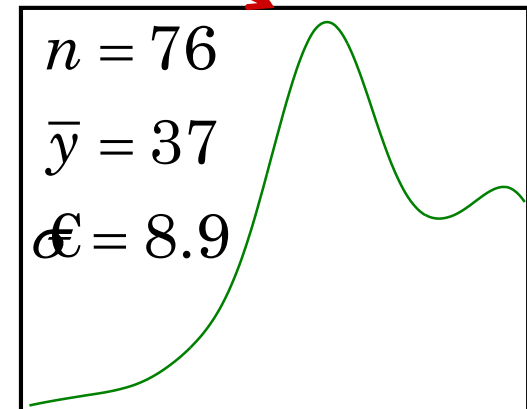
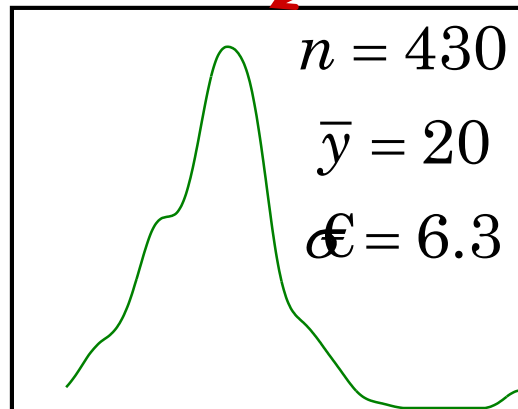
# Уменьшение вариации



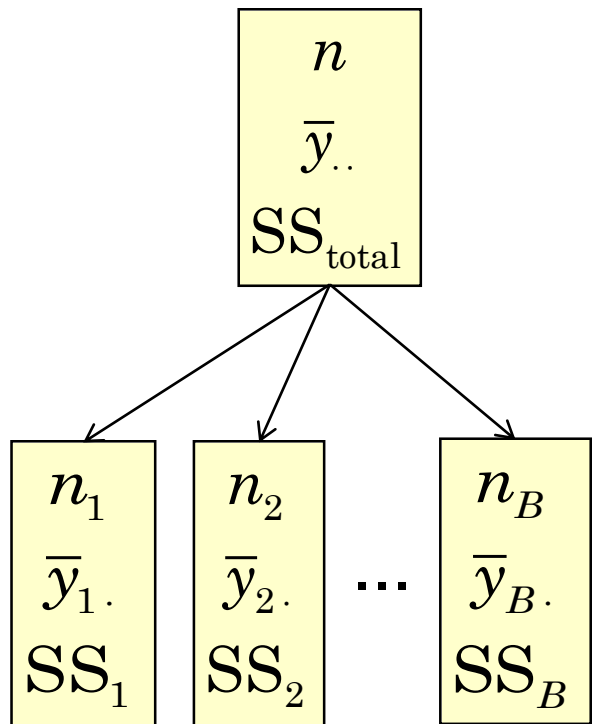
yes

$RM < 6.94$

no



# One-Way ANOVA



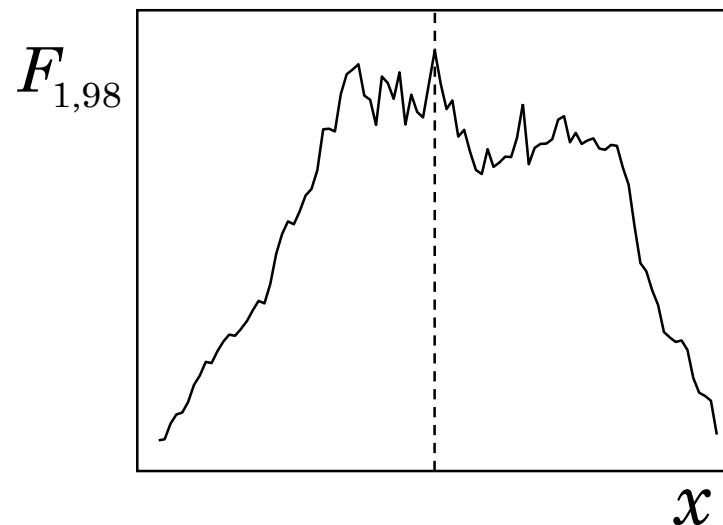
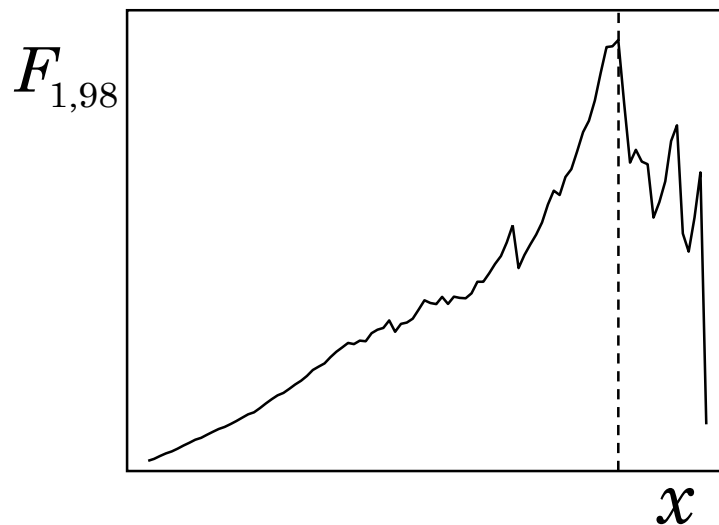
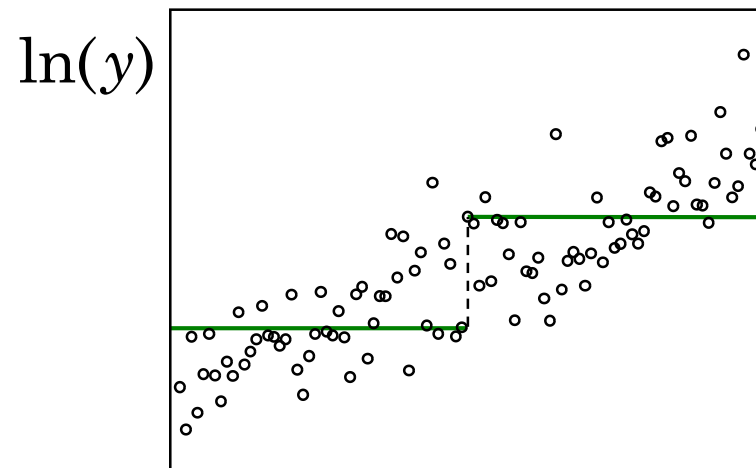
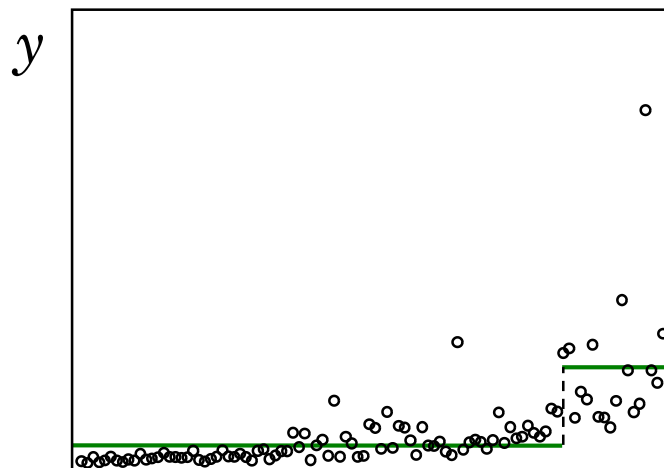
$$F = \left( \frac{SS_{\text{between}}}{SS_{\text{within}}} \right) \left( \frac{n - B}{B - 1} \right) \sim F_{B-1, n-B}$$

$$SS_{\text{between}} = \sum_{i=1}^B n_i (\bar{y}_{i.} - \bar{y}_{..})^2$$

$$SS_{\text{within}} = \sum_{i=1}^B SS_i = \sum_{i=1}^B \sum_{j=1}^{n_i} (y_{ij} - \bar{y}_{i.})^2$$

$$SS_{\text{total}} = \sum_{i=1}^B \sum_{j=1}^{n_i} (y_{ij} - \bar{y}_{..})^2$$

# Гетероскедастичность

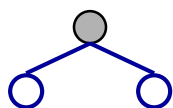


# Переобучение деревьев решений

- Описанный выше процесс может давать хорошие прогнозы для обучающего набора, но может привести к *переобучению*, что приведет к ухудшению качества на тестовом наборе.
- Меньшее по размеру дерево с меньшим количеством разбиений (то есть меньшим количеством областей  $R_1, \dots, R_J$ ) может привести к снижению дисперсии и лучшей интерпретируемости.
- Одной из возможных альтернатив является ранняя остановка построения дерева, определяемая параметрами:
  - Максимальная глубина дерева
  - Минимальное число наблюдений в листе
  - Порог на p-value

# Рост дерева

| $-\log_{10}(P)$ | $m$ | $-\log_{10}(mP)$ | $d$ | $-\log_{10}(2^d mP)$ |
|-----------------|-----|------------------|-----|----------------------|
|-----------------|-----|------------------|-----|----------------------|



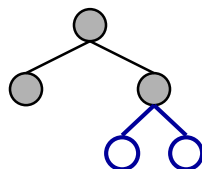
26.7

53

24.9

0

24.9



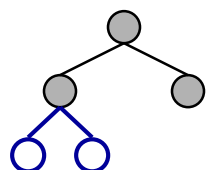
3.12

14

1.97

1

1.67



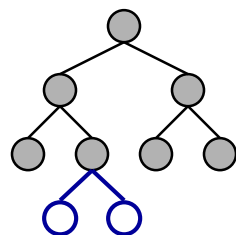
1.63

39

.039

1

-.26



2.40

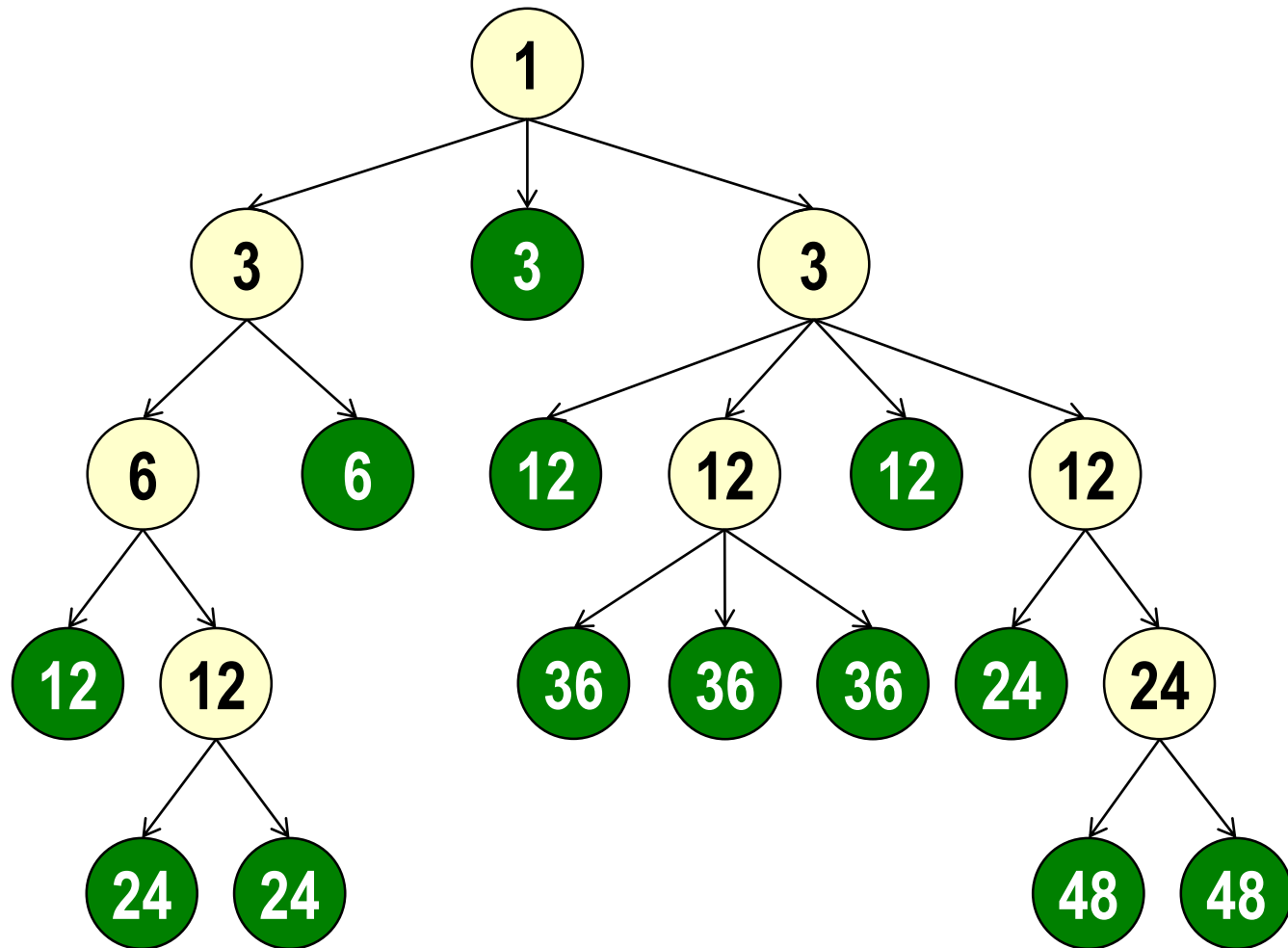
11

1.36

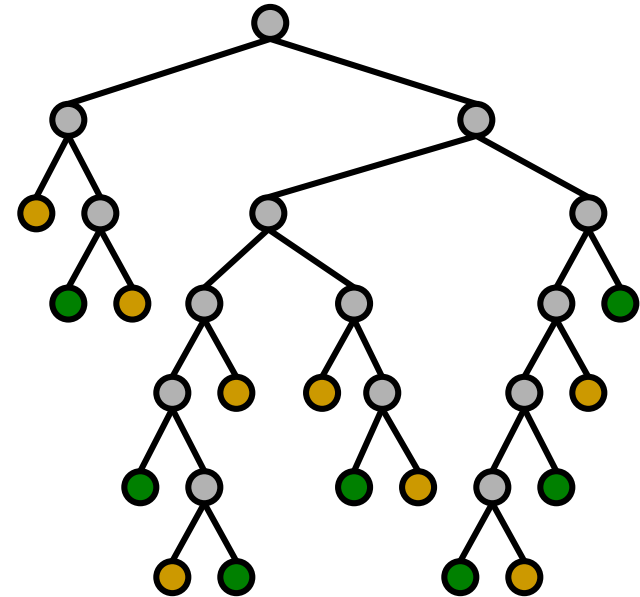
2

.76

# Множитель глубины

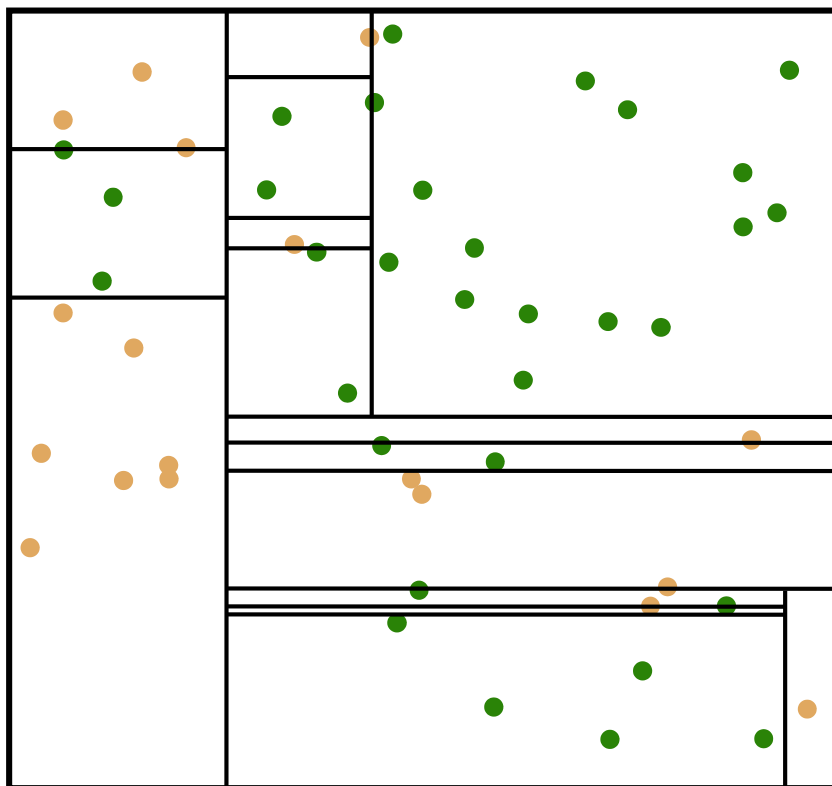


The diagram illustrates a hierarchical tree structure, likely representing a decision tree or a classification model. The tree is rooted at the top left and branches outwards. The nodes are represented by circles, and the edges are represented by lines. The nodes are colored either green or orange, indicating different classes or states. The tree structure is complex, with multiple levels of branching and a large number of nodes. The overall shape of the tree is roughly rectangular, with the root at the top left and the leaves extending towards the bottom right.

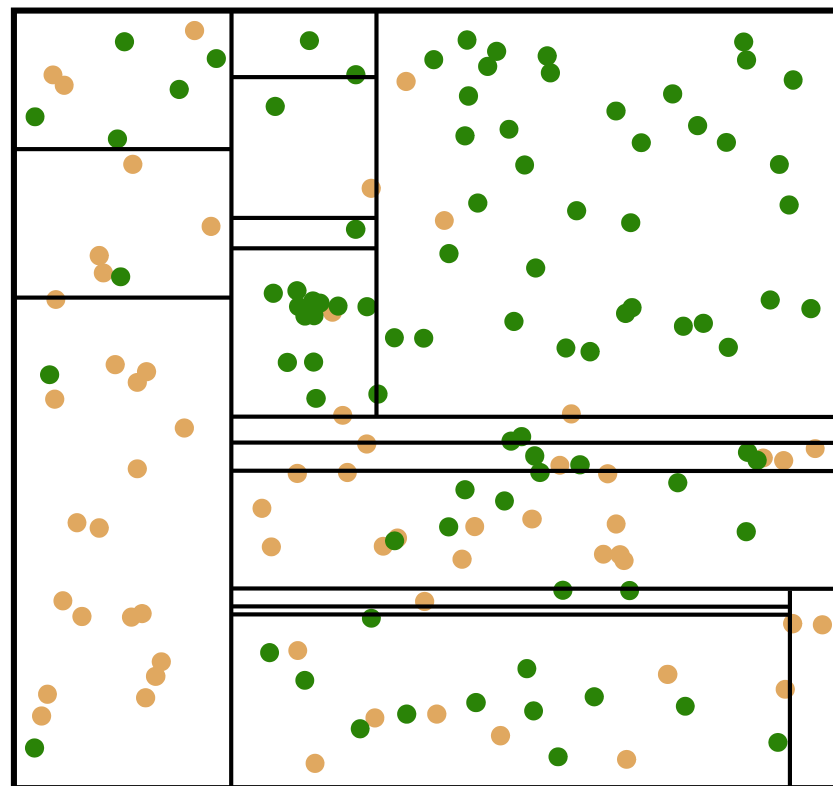


# Переобучение

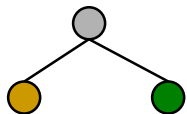
Тренировочный набор



Новые данные

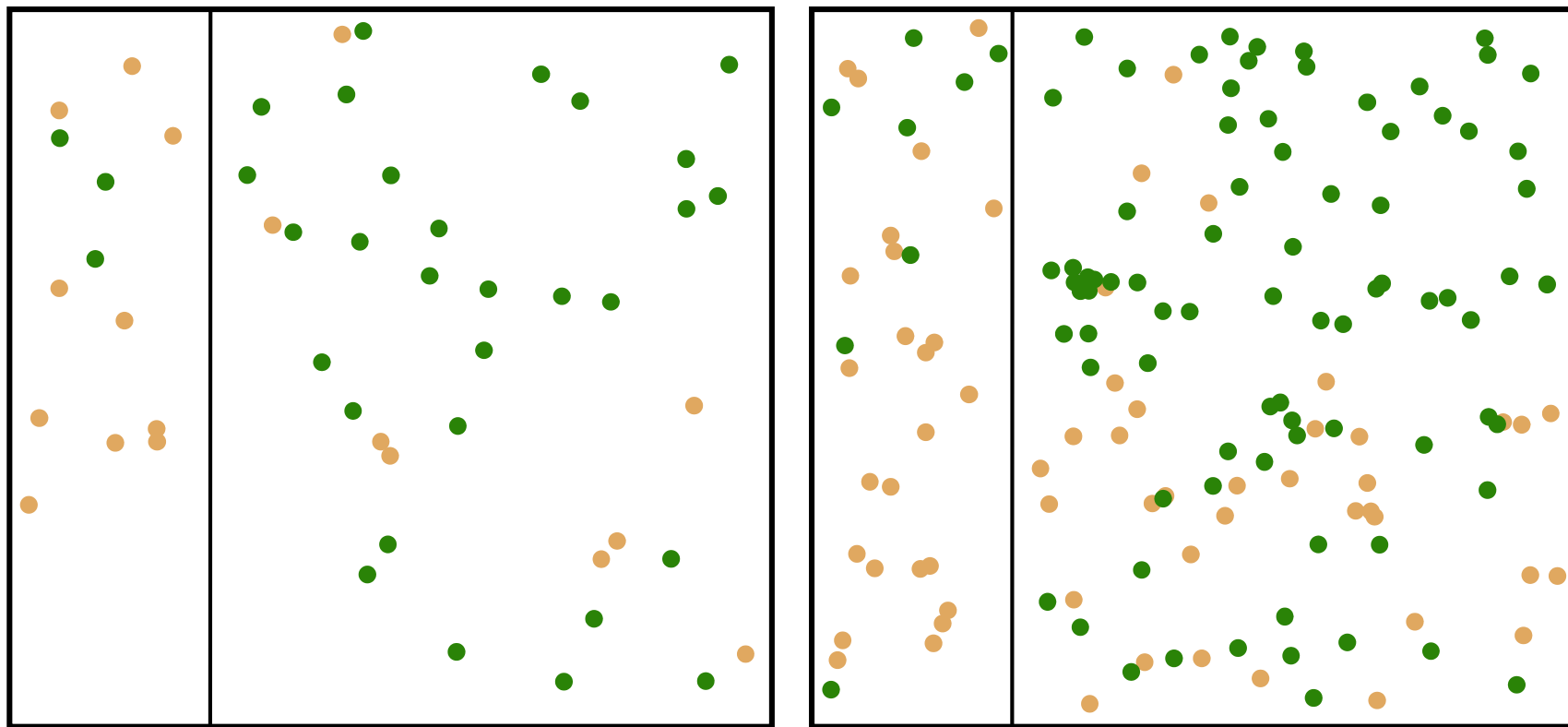


# Недообучение



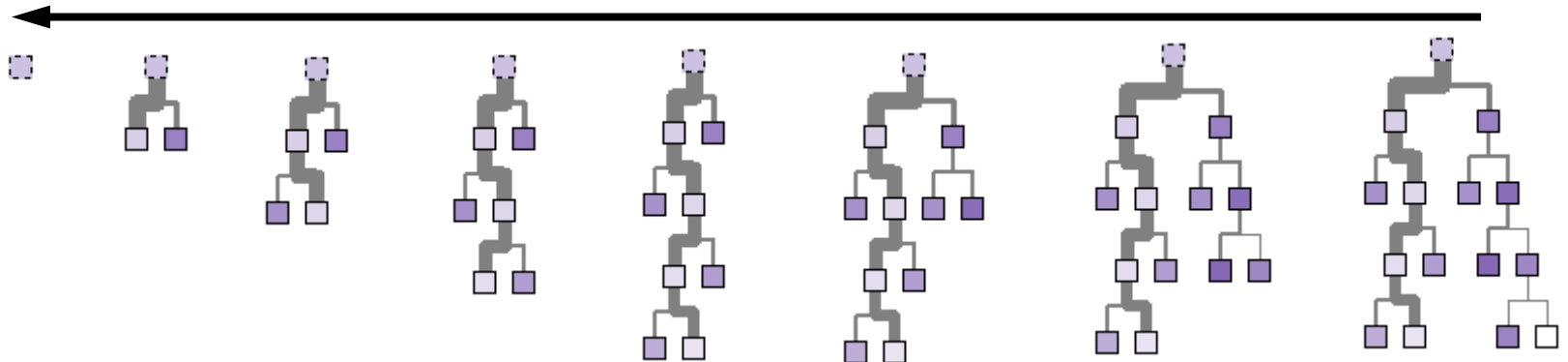
Тренировочный набор

Новые данные



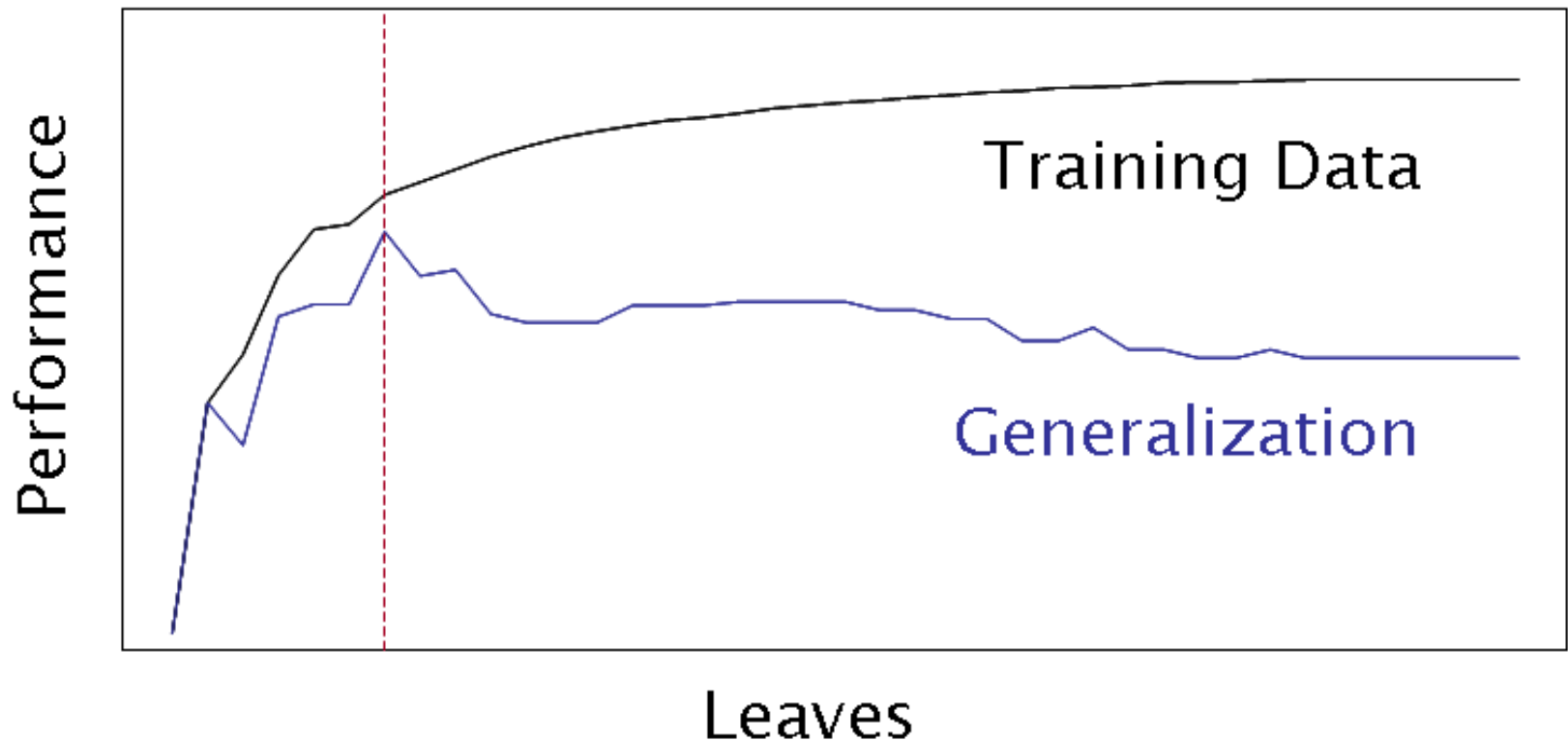
# Обрубание дерева

- Лучшей стратегией является построение большого дерева  $T_0$ , а затем выполнение *отсечения* для получения *поддерева*
- Для этого используется подход *сокращения сложности*, также известный как *удаление самых слабых связей*
- Используем валидационный набор или кросс валидацию для выбора оптимального размера дерева.
- Алгоритм: строим максимальное дерево и последовательно обрубаем ветки с отбором

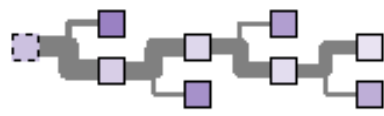
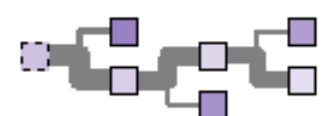
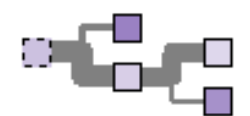
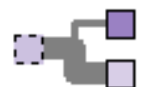


# Обрубание ветвей

Выбираем лучшее на валидационном наборе.



# Критерий выбора поддерева

| <u>Leaves</u> | <u>Accuracy</u> | <u>Profit</u> | <u>ASE</u> |                                                                                       |
|---------------|-----------------|---------------|------------|---------------------------------------------------------------------------------------|
| 5             | .90/.88         | .51/.43       | .17/.15    |    |
| 4             | .90/.88         | .51/.40       | .18/.14    |    |
| 3             | .89/.91         | .49/.44       | .19/.16    |    |
| 2             | .88/.91         | .49/.44       | .20/.16    |  |
| 1             | .59/.64         | .04/ -.1      | .48/.46    |  |

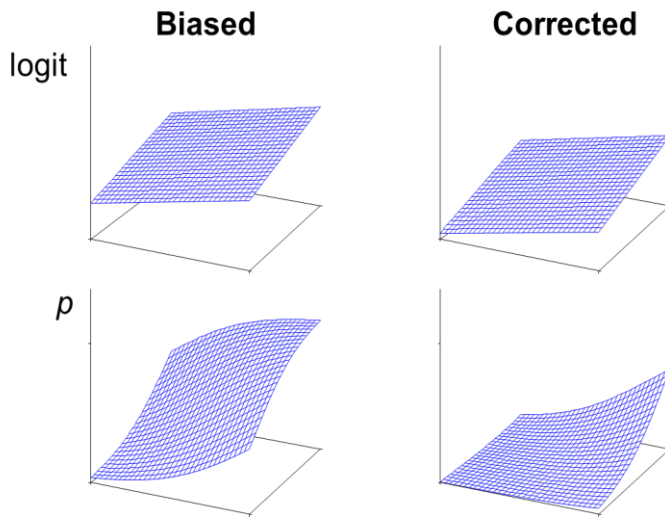
# Оценка точности

|                |   | Решение/Действие |                 | Скорректированная                    |                                      |
|----------------|---|------------------|-----------------|--------------------------------------|--------------------------------------|
|                |   | 0                | 1               |                                      |                                      |
| Реальный класс | 0 | $n_{\text{TN}}$  | $n_{\text{FP}}$ | $\frac{\pi_0}{\rho_0} n_{\text{TN}}$ | $\frac{\pi_0}{\rho_0} n_{\text{FP}}$ |
|                | 1 | $n_{\text{FN}}$  | $n_{\text{TP}}$ | $\frac{\pi_1}{\rho_1} n_{\text{FN}}$ | $\frac{\pi_1}{\rho_1} n_{\text{TP}}$ |

$$\text{Accuracy} = \frac{1}{n} \left( \frac{\pi_0}{\rho_0} n_{\text{TN}} + \frac{\pi_1}{\rho_1} n_{\text{TP}} \right)$$

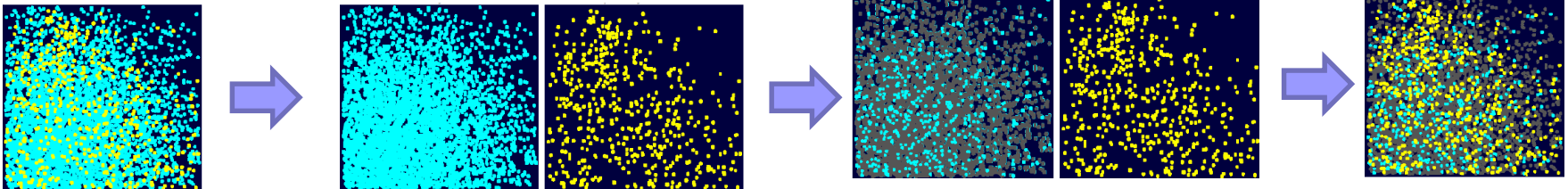
# «Балансировка» выборки (oversampling)

- Порог отсечения для логистической функции:

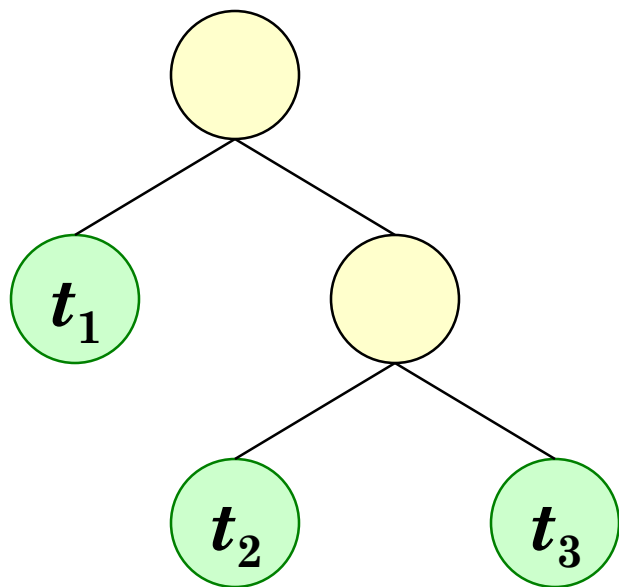


- «Балансировка»:

| Property                   | Value       |
|----------------------------|-------------|
| <b>General</b>             |             |
| Node ID                    | Smpl        |
| Imported Data              | ...         |
| Exported Data              | ...         |
| Notes                      | ...         |
| <b>Train</b>               |             |
| Variables                  | ...         |
| Output Type                | Data        |
| Sample Method              | Default     |
| Random Seed                | 12345       |
| <b>Size</b>                |             |
| Type                       | Percentage  |
| Observations               | .           |
| Percentage                 | 10.0        |
| Alpha                      | 0.01        |
| PValue                     | 0.01        |
| Cluster Method             | Random      |
| <b>Stratified</b>          |             |
| Criterion                  | Level Based |
| Ignore Small Strata        | No          |
| Minimum Strata Size        | 5           |
| <b>Level Based Options</b> |             |
| Level Selection            | Event       |
| Level Proportion           | 100.0       |
| Sample Proportion          | 50.0        |
| <b>Oversampling</b>        |             |
| Adjust Frequency           | No          |
| Based on Count             | No          |

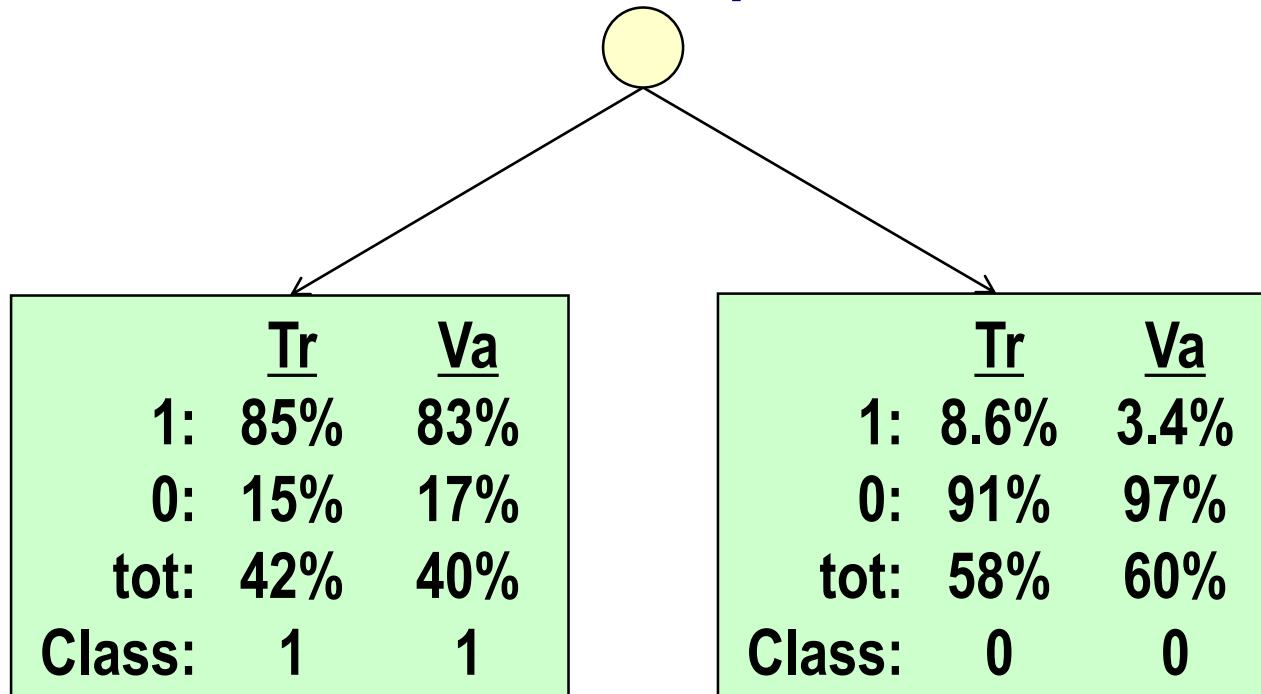


# Точность дерева



$$\text{Accuracy} = \frac{1}{n} \left( n(t_1) \text{acc}(t_1) + n(t_2) \text{acc}(t_2) + n(t_3) \text{acc}(t_3) \right)$$

# Максимизация точности



$$\text{Training Accuracy} = (.42)(.85) + (.58)(.91) = .88$$

$$\text{Validation Accuracy} = (.40)(.83) + (.60)(.97) = .91$$

# Матрица выигрыша

|                |   | Решение              |                      |
|----------------|---|----------------------|----------------------|
|                |   | 0                    | 1                    |
| Реальный класс | 0 | $\delta_{\text{TN}}$ | $\delta_{\text{FP}}$ |
|                | 1 | $\delta_{\text{FN}}$ | $\delta_{\text{TP}}$ |

Правило Байеса:

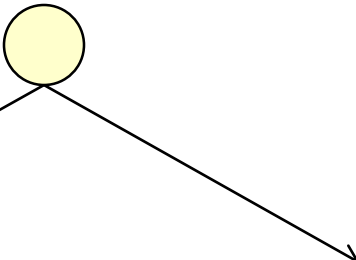
Решение 1 если

$$P > \frac{1}{1 + \left( \frac{\delta_{\text{TP}} - \delta_{\text{FN}}}{\delta_{\text{TN}} - \delta_{\text{FP}}} \right)}$$

Profit =

$$\frac{1}{n} \left( \delta_{\text{TN}} \frac{\pi_0}{\rho_0} n_{\text{TN}} + \delta_{\text{FP}} \frac{\pi_0}{\rho_0} n_{\text{FP}} + \delta_{\text{FN}} \frac{\pi_1}{\rho_1} n_{\text{FN}} + \delta_{\text{TP}} \frac{\pi_1}{\rho_1} n_{\text{TP}} \right)$$

# Максимизация выигрыша



|        | <u>Tr</u> | <u>Va</u> |
|--------|-----------|-----------|
| 1:     | 85%       | 83%       |
| 0:     | 15%       | 18%       |
| tot:   | 42%       | 40%       |
| P1:    | 1.18      | 1.11      |
| P0:    | 0         | 0         |
| Class: | 1         | 1         |

|        | <u>Tr</u> | <u>Va</u> |
|--------|-----------|-----------|
| 1:     | 8.6%      | 3.4%      |
| 0:     | 91%       | 97%       |
| tot:   | 58%       | 60%       |
| P1:    | -.78      | -.91      |
| P0:    | 0         | 0         |
| Class: | 0         | 0         |

|            |   |         |   |
|------------|---|---------|---|
|            |   | прогноз |   |
|            |   | 1       | 0 |
| реальность | 1 | 1.56    | 0 |
|            | 0 | -1      | 0 |

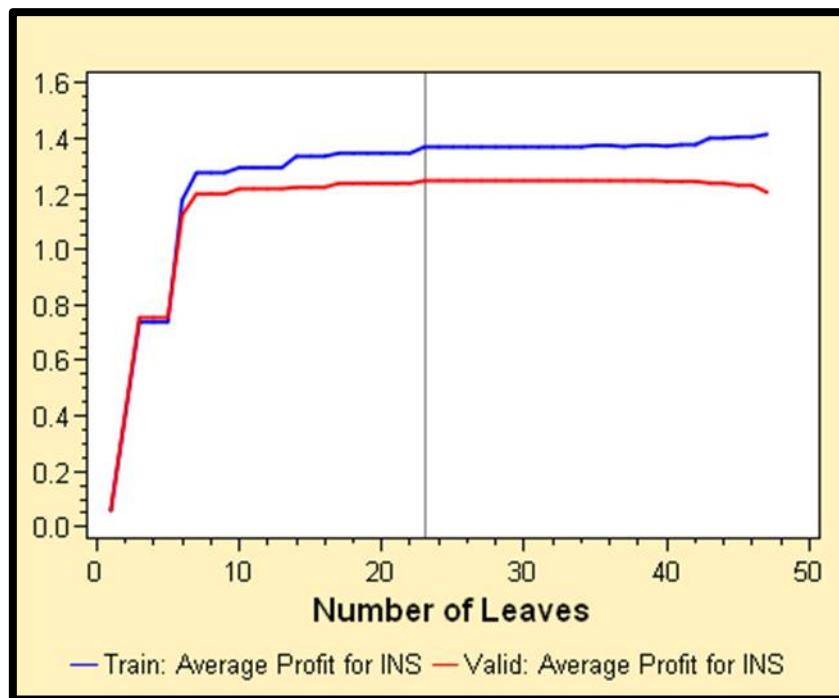
Матрица выигрыша

$$\text{Training Profit} = (.42)(1.18) + (.58)(0) = .50$$

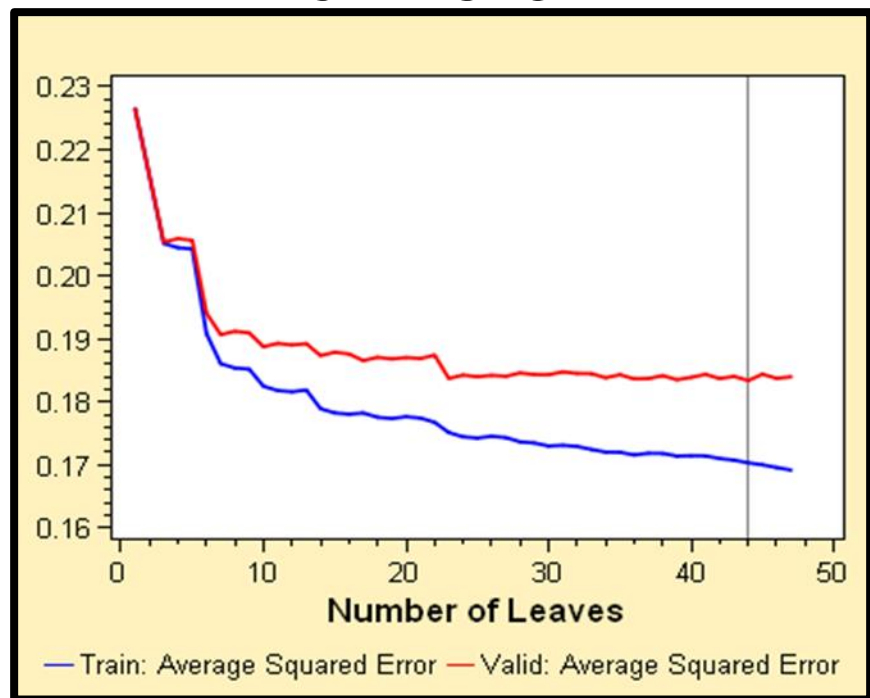
$$\text{Validation Profit} = (.40)(1.11) + (.60)(0) = .44$$

# Вероятностное дерево

Выигрыш



Среднеквадратичная  
ошибка



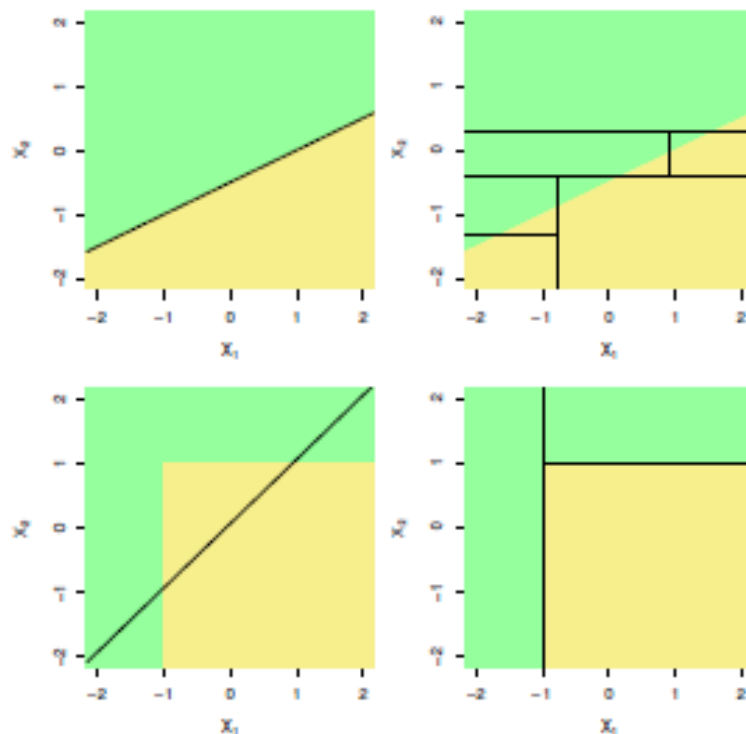
$$ASE_{nominal} = \frac{1}{n} \sum_{i=1}^n \sum_{j=1}^k (y_i - p_k(x_i))^2$$

$$ASE_{interval} = \frac{1}{n} \sum_{i=1}^n (y_i - f(x_i))^2$$

# Особенности популярных алгоритмов построения деревьев решений

| Свойства                   | CHAID (Kass)                       | CART (Breiman)                    | C5 (Quinlan)                 |
|----------------------------|------------------------------------|-----------------------------------|------------------------------|
| Критерий для числ. отклика | Фишер                              | Вариация                          | нет                          |
| Критерий для симв. отклика | Хи-квадрат                         | Джини                             | Энтропия                     |
| Работа с пропусками        | Отдельная ветвь                    | Подстановка                       | Пропорция по ветвям (игнор)  |
| Особенности                | Корректировка Бонферрони и глубина | Линейные комбинации при разбиении | Алгоритм «сокращения» правил |
| Обрубание вервей           | НЕТ                                | По точности                       | По точности                  |

# Деревья решений vs линейные модели



Верхний ряд: истинная линейная граница; Нижний ряд: истинная нелинейная граница.

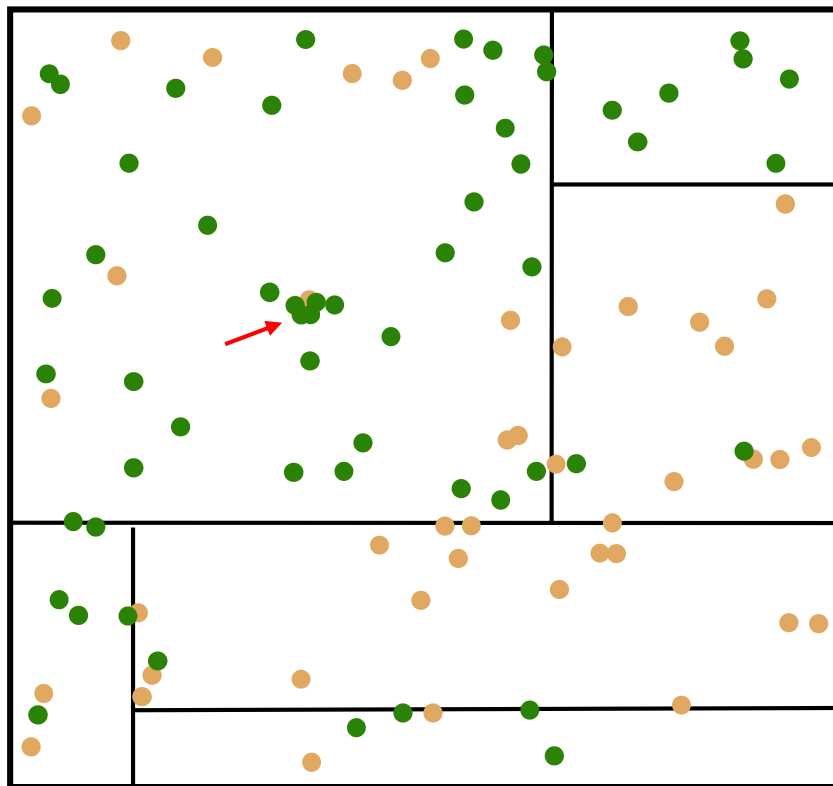
Левый столбец: линейная модель; Правый столбец: модель на основе деревьев решений

# Преимущества и недостатки деревьев решений

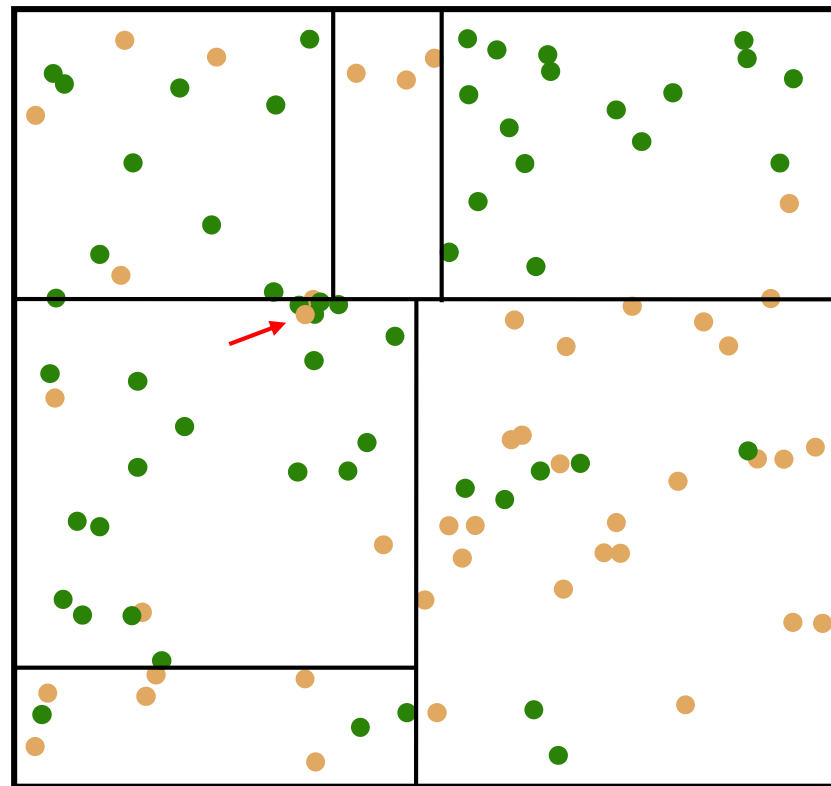
- Деревья имеют очень понятную интерпретацию. Даже проще, чем линейная регрессия!
- Существует мнение, что деревья решений отражают процесс принятия решений людьми лучше, чем подходы для решения задач регрессии и классификации, рассмотренные в предыдущих главах.
- Деревья можно наглядно отобразить графически и легко интерпретировать даже не специалистам (особенно, для небольших деревьев).
- Деревья могут легко обрабатывать качественные переменные и пропуски без необходимости создания фиктивных переменных.
- К сожалению, деревья обычно не дают такую же точность прогнозирования, как некоторые другие подходы для решения задач регрессии и классификации.

# Нестабильность модели

Один новый пример



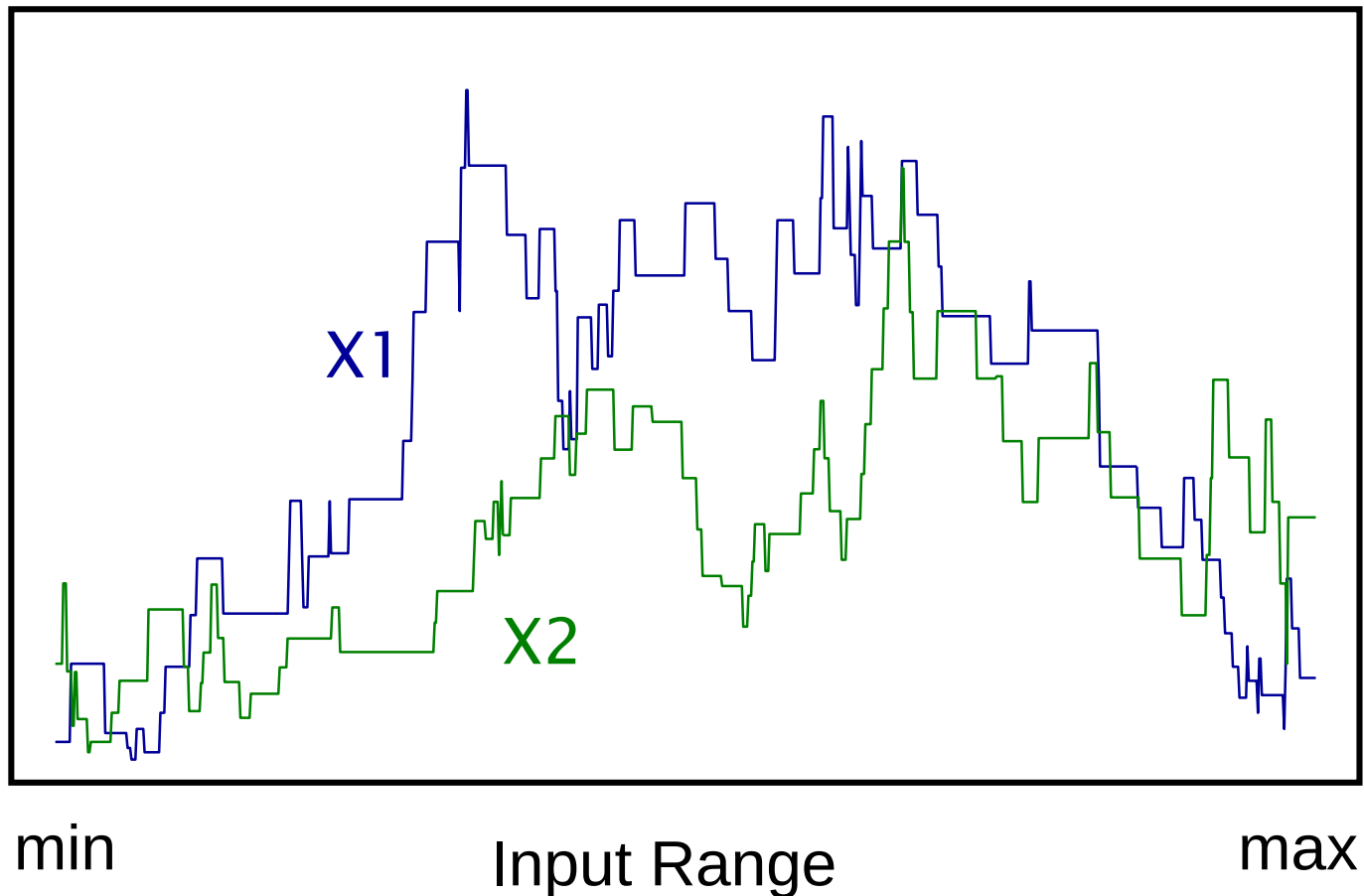
Точность = 81%



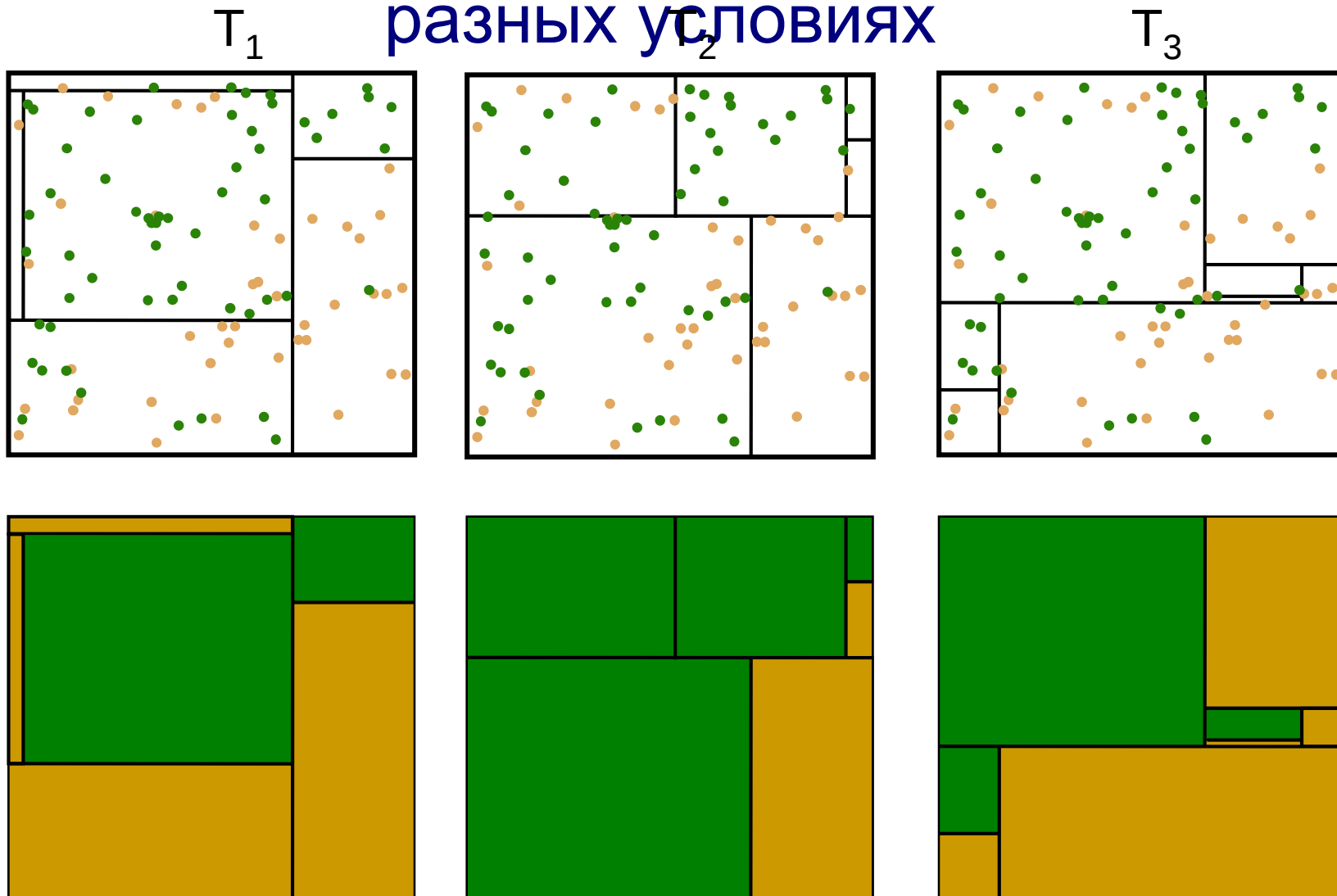
Точность = 80%

# Альтернативные точки разбиения

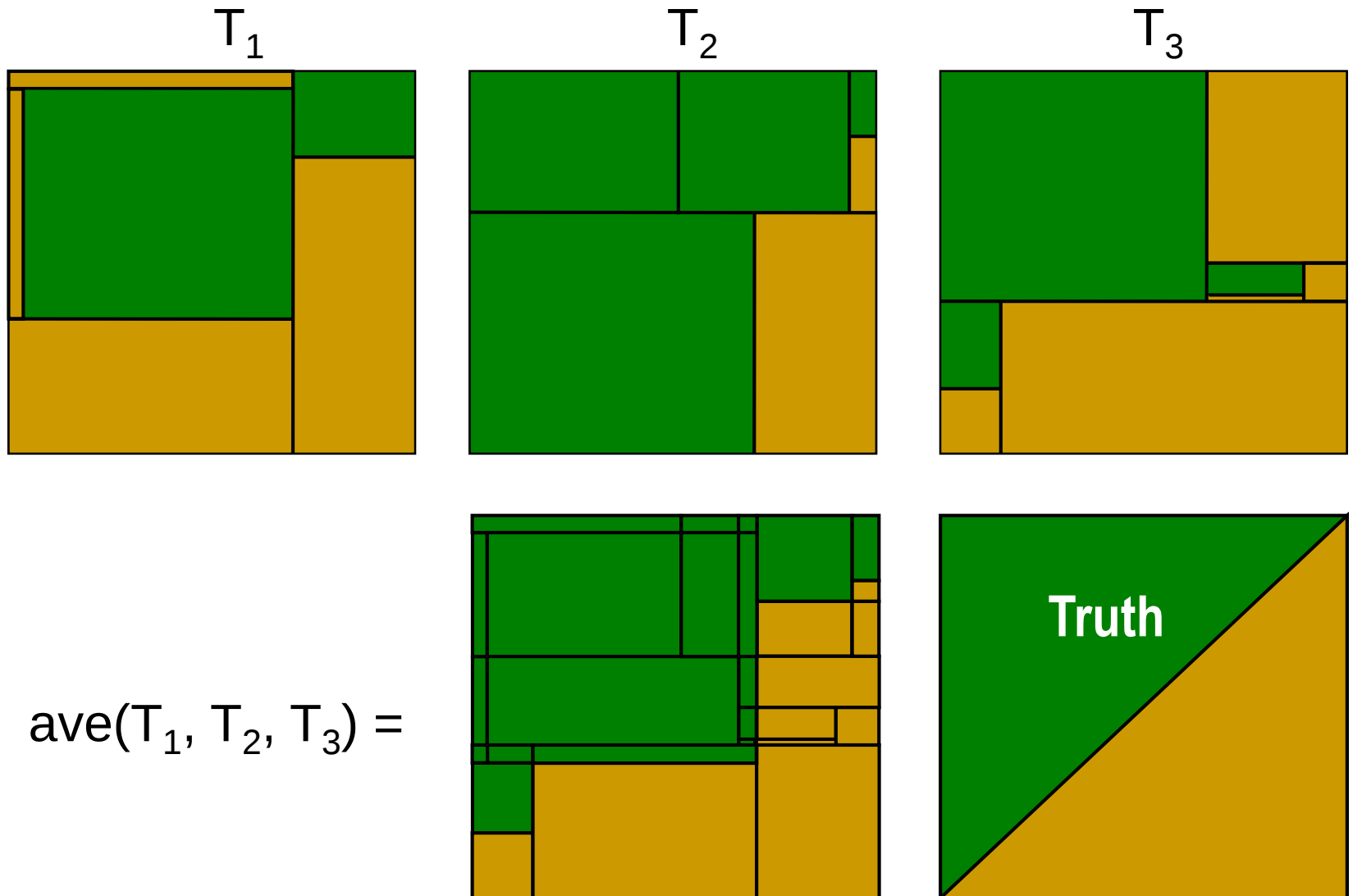
Logworth



# НЕСКОЛЬКО моделей построенных в разных условиях

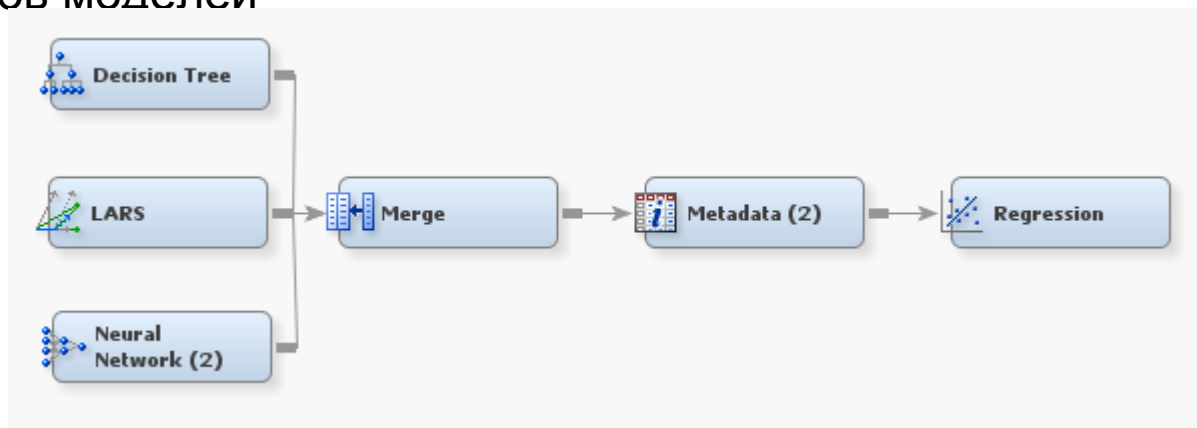
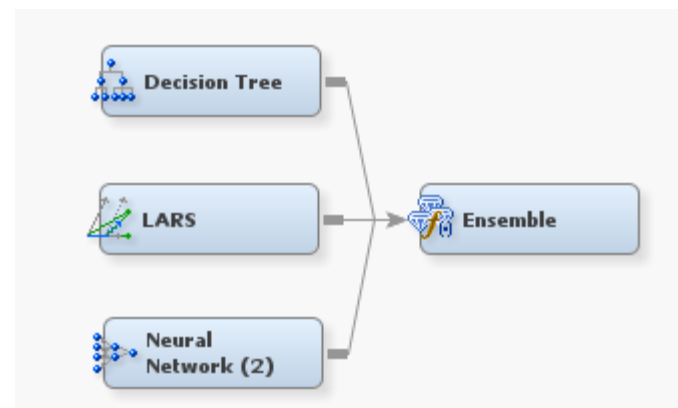


# Ансамбль = Комбинация моделей

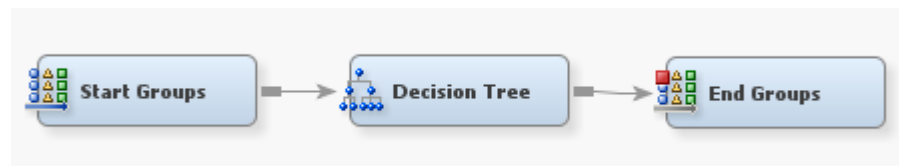


# Типы ансамблей

- Простое «голосование»
  - Усреднение отклика, максимум отклика, «пропорция» голосов
- Взвешенное «голосование»
  - Строится новая простая модель (например регрессия или простая нейронная сеть) на основе откликов моделей ансамбля



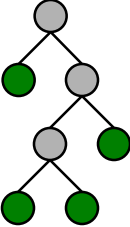
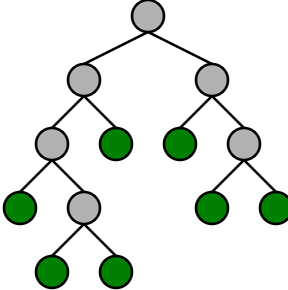
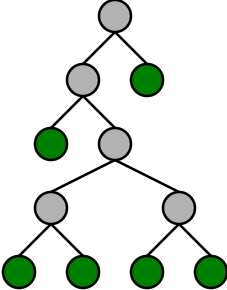
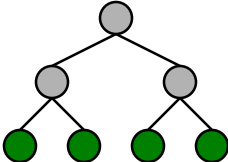
- «Комбинации»
  - Bagging - «усреднение» прогноза на разных выборках
  - Boosting - «усиление» прогноза на разных выборках



# Bagging

|             | k=1         | k=2         | k=3         | k=4 ...     |
|-------------|-------------|-------------|-------------|-------------|
| <u>case</u> | <u>freq</u> | <u>freq</u> | <u>freq</u> | <u>freq</u> |
| 1           | 1           | 0           | 3           | 1           |
| 2           | 0           | 1           | 1           | 1           |
| 3           | 2           | 0           | 0           | 2           |
| 4           | 0           | 2           | 2           | 0           |
| 5           | 2           | 2           | 0           | 1           |
| 6           | 1           | 1           | 0           | 1           |

# Arc-x4

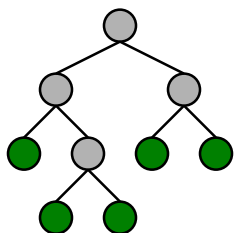
k=1

k=2

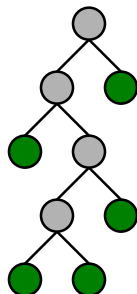
k=3

k=4 ...

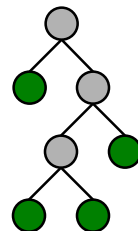
| <u>case</u> | <u>freq</u> | <u>m</u> |
|-------------|-------------|----------|
| 1           | 1           | 1        |
| 2           | 1           | 0        |
| 3           | 1           | 1        |
| 4           | 1           | 0        |
| 5           | 1           | 0        |
| 6           | 1           | 0        |



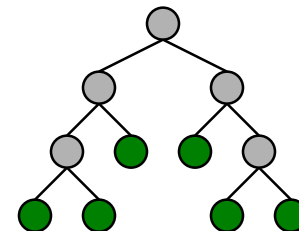
| <u>freq</u> | <u>m</u> |
|-------------|----------|
| 1.5         | 1        |
| .75         | 0        |
| 1.5         | 2        |
| .75         | 1        |
| .75         | 0        |
| .75         | 0        |



| <u>freq</u> | <u>m</u> |
|-------------|----------|
| .5          | 2        |
| .25         | 0        |
| 4.25        | 3        |
| .5          | 1        |
| .25         | 0        |
| .25         | 1        |

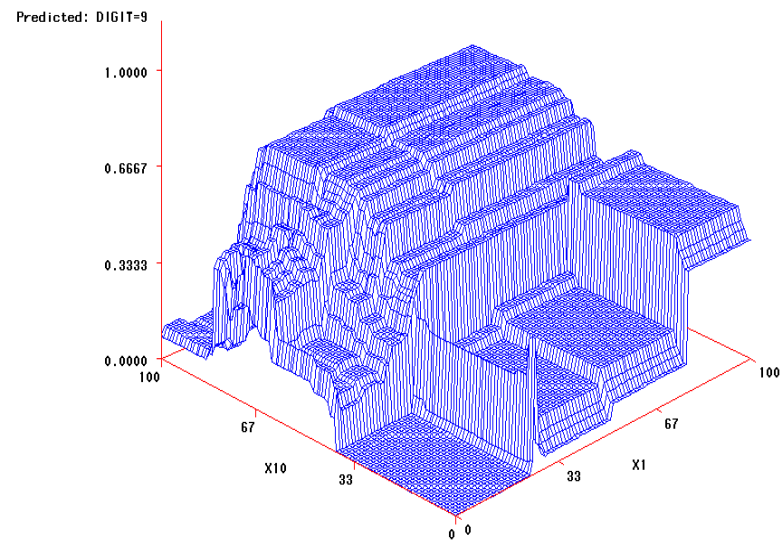
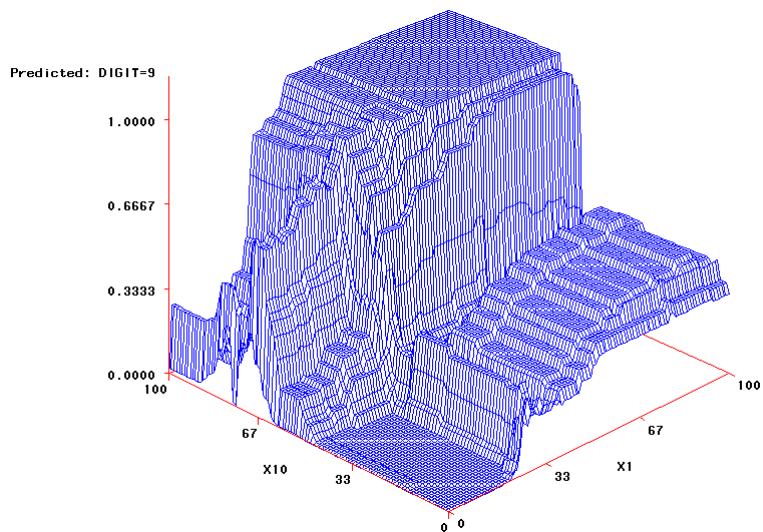
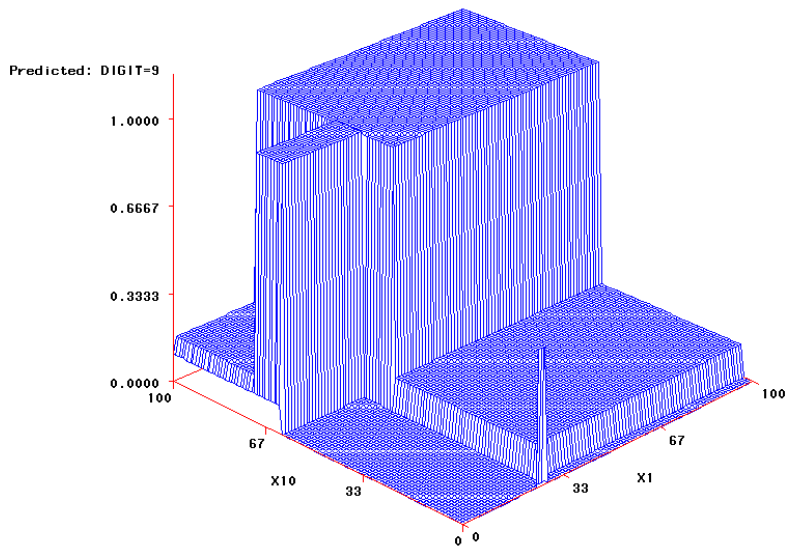


| <u>freq</u> |
|-------------|
| .97         |
| .06         |
| 4.69        |
| .11         |
| .06         |
| .11         |



$$p(i) = \frac{1 + m(i)^4}{\sum (1 + m(i)^4)}$$

# Обычное, Bagged и Boosted Деревья



# Логистическая регрессия

Почему нельзя моделировать вероятность отклика  $p$  как непрерывный отклик с помощью линейной регрессии?

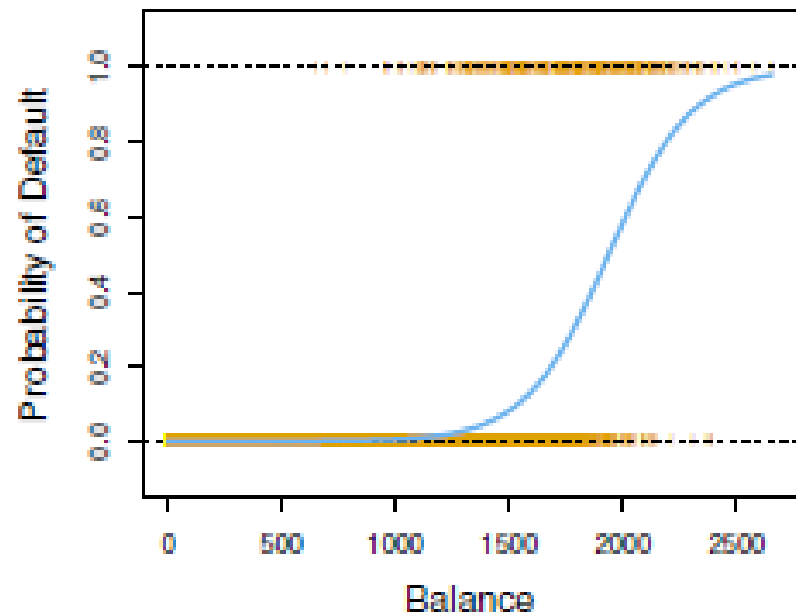
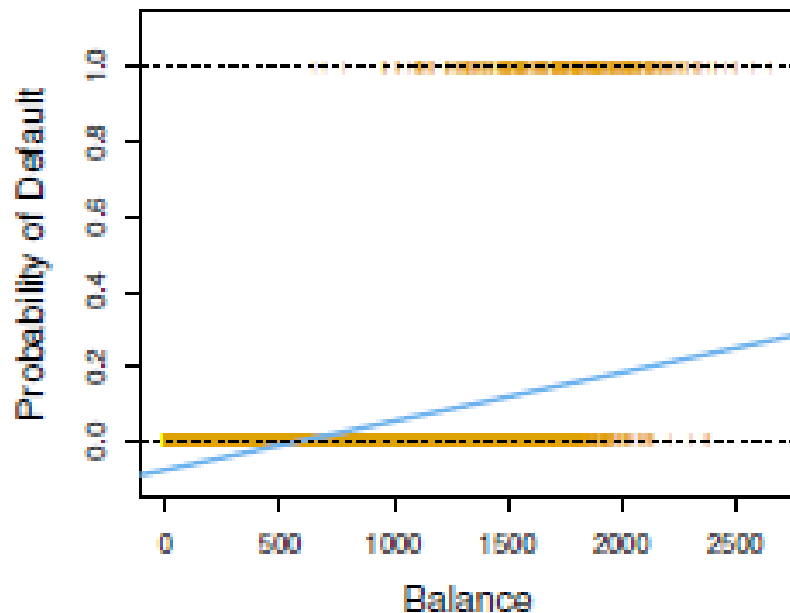
OLS Reg:  $Y_i = \beta_0 + \beta_1 X_{1i} + \varepsilon_i$

- Если целевая переменная категориальная, как представить ее в виде числовой?
- Если целевая закодирована (1=Yes and 0=No) а результат модели 0.5 или 1.1 или -0.4, что это означает?
- Если переменная имеет только два значения (или несколько), имеет ли смысл требовать постоянства дисперсии или нормальности ошибок?

Linear Prob. Model:  $p_i = \beta_0 + \beta_1 X_{1i}$

- Вероятность ограничена, а линейная функция принимает любые значения.
- Принимая во внимание ограниченность вероятности, можно ли предполагать линейную связь между  $X$  и  $p$ ?
- Можно ли предполагать ошибку с постоянной дисперсией?
- Что такое наблюдаемая вероятность для конкретного наблюдения? 0 и 1?

# Линейная vs логистическая регрессия



Логистическая регрессия гарантирует, что оценка  $p(X)$  находится в диапазоне между 0 и 1.

# Логистическая регрессия

Уравнение логистической регрессии: **Функция связи** (логит) и обратная ей (логистическая):

вероятность

$$\text{logit}(p_i) = \beta_0 + \beta_1 x_{1i} + \dots + \beta_k x_{ki}$$

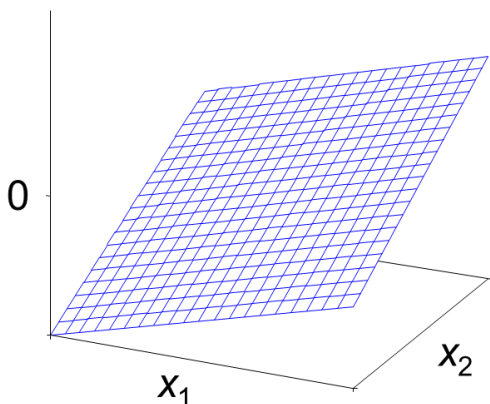
параметр      предиктор

$$\text{logit}(p_i) = \ln\left(\frac{p_i}{1-p_i}\right) = \eta$$

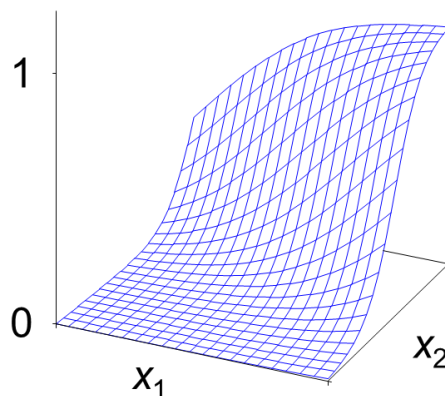
$$\Leftrightarrow p_i = \frac{1}{1+e^{-\eta}}$$

Основное предположение линейной логистической регрессии (линейная зависимость логита от предикторов):

$\text{logit}(p)$

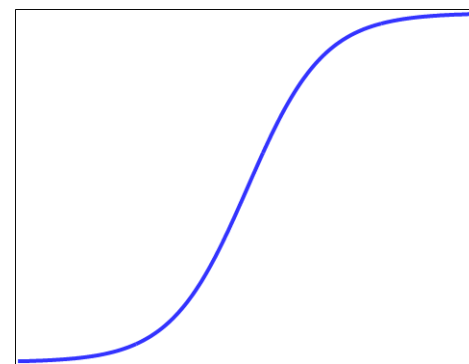


$p$



$p_i = 1$

$p_i = 0$



меньше  $\leftarrow \eta \rightarrow$  больше

Ограничивает значение отклика

# Максимальное правдоподобие

Максимальное правдоподобие для оценки параметров.

$$\ell(\beta_0, \beta) = \prod_{i: y_i=1} p(x_i) \prod_{i: y_i=0} (1 - p(x_i)).$$

Это правдоподобие дает вероятность наблюдаемых нулей и единиц в данных. Мы выбираем параметры модели, чтобы максимизировать вероятность наблюдаемых данных.

$$\log \left( \frac{p(X)}{1 - p(X)} \right) = \beta_0 + \beta_1 X_1 + \dots + \beta_p X_p$$

$$p(X) = \frac{e^{\beta_0 + \beta_1 X_1 + \dots + \beta_p X_p}}{1 + e^{\beta_0 + \beta_1 X_1 + \dots + \beta_p X_p}}$$

|               | Coefficient | Std. Error | Z-statistic | P-value  |
|---------------|-------------|------------|-------------|----------|
| Intercept     | -10.8690    | 0.4923     | -22.08      | < 0.0001 |
| balance       | 0.0057      | 0.0002     | 24.74       | < 0.0001 |
| income        | 0.0030      | 0.0082     | 0.37        | 0.7115   |
| student [Yes] | -0.6468     | 0.2362     | -2.74       | 0.0062   |

# Отношение шансов

- Показывает как изменится отношение шансов при изменении  $i$ -ой переменной на 1 unit (равно exp от коэф.)

$$\text{logit}(\hat{p}) = \log(\text{odds}) = \beta_0 + \beta_i * x_i + \sum_{j \neq i} \beta_j * x_j$$

$$\text{odds} = \exp(\beta_0 + \beta_i * x_i + \sum_{j \neq i} \beta_j * x_j)$$

$$\text{logit}(\hat{p}') = \log(\text{odds}) = \beta_0 + \beta_i * (x_i + 1) + \sum_{j \neq i} \beta_j * x_j$$

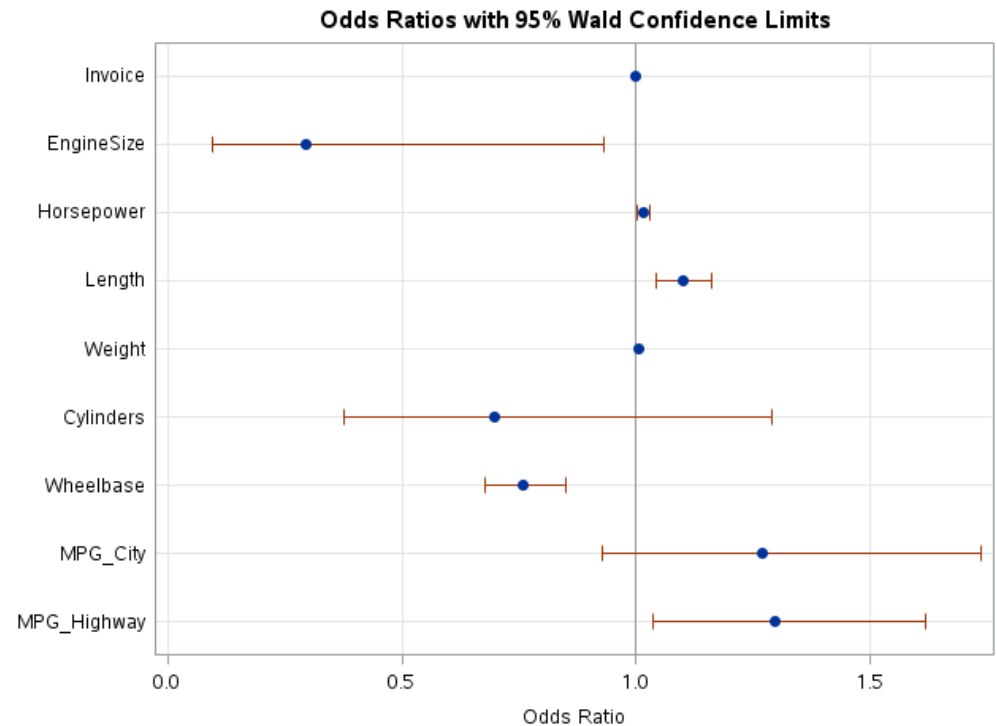
$$\text{odds}' = \exp(\beta_0 + \beta_i * (x_i + 1) + \sum_{j \neq i} \beta_j * x_j)$$

$$\text{odds ratio} = \text{odds}' / \text{odds} = \exp(\beta_i)$$

Больше 1 – отношение шансов увеличивается, если меньше, то уменьшается

# Отношение шансов (пример)

| Odds Ratio Estimates |                |                            |       |
|----------------------|----------------|----------------------------|-------|
| Effect               | Point Estimate | 95% Wald Confidence Limits |       |
| Invoice              | 1.000          | 1.000                      | 1.000 |
| EngineSize           | 0.295          | 0.094                      | 0.931 |
| Horsepower           | 1.016          | 1.003                      | 1.029 |
| Length               | 1.100          | 1.044                      | 1.160 |
| Weight               | 1.005          | 1.004                      | 1.007 |
| Cylinders            | 0.696          | 0.376                      | 1.289 |
| Wheelbase            | 0.757          | 0.676                      | 0.849 |
| MPG_City             | 1.270          | 0.929                      | 1.736 |
| MPG_Highway          | 1.295          | 1.036                      | 1.618 |



# Категориальные предикторы

## □ Схемы кодировки:

- Effect coding (относительно «среднего»)

| <u>CLASS</u> | <u>Value</u> | <u>Label</u>  | <u>1</u> | <u>2</u> |
|--------------|--------------|---------------|----------|----------|
| IncLevel     | 1            | Low Income    | 1        | 0        |
|              | 2            | Medium Income | 0        | 1        |
|              | 3            | High Income   | -1       | -1       |

- Reference coding (относительно «базового»)

| <u>CLASS</u> | <u>Value</u> | <u>Label</u>  | <u>1</u> | <u>2</u> |
|--------------|--------------|---------------|----------|----------|
| IncLevel     | 1            | Low Income    | 1        | 0        |
|              | 2            | Medium Income | 0        | 1        |
|              | 3            | High Income   | 0        | 0        |

# Effect Coding: Пример

$$\text{logit}(p) = \beta_0 + \beta_1 * D_{\text{Low income}} + \beta_2 * D_{\text{Medium income}}$$

$\beta_0$  = Средний логит по всем категориям

$\beta_1$  = Разница между логитом для Low income и средним логитом

$\beta_2$  = разница между Medium income и средним логитом

| Analysis of Maximum Likelihood Estimates |   |    |          |                |                 |            |
|------------------------------------------|---|----|----------|----------------|-----------------|------------|
| Parameter                                |   | DF | Estimate | Standard Error | Wald Chi-Square | Pr > ChiSq |
| Intercept                                |   | 1  | -0.5363  | 0.1015         | 27.9143         | <.0001     |
| IncLevel                                 | 1 | 1  | -0.2259  | 0.1481         | 2.3247          | 0.1273     |
| IncLevel                                 | 2 | 1  | -0.2200  | 0.1447         | 2.3111          | 0.1285     |

# Reference Coding: Пример

$$\text{logit}(p) = \beta_0 + \beta_1 * D_{\text{Low income}} + \beta_2 * D_{\text{Medium income}}$$

$\beta_0$  = Логит для High

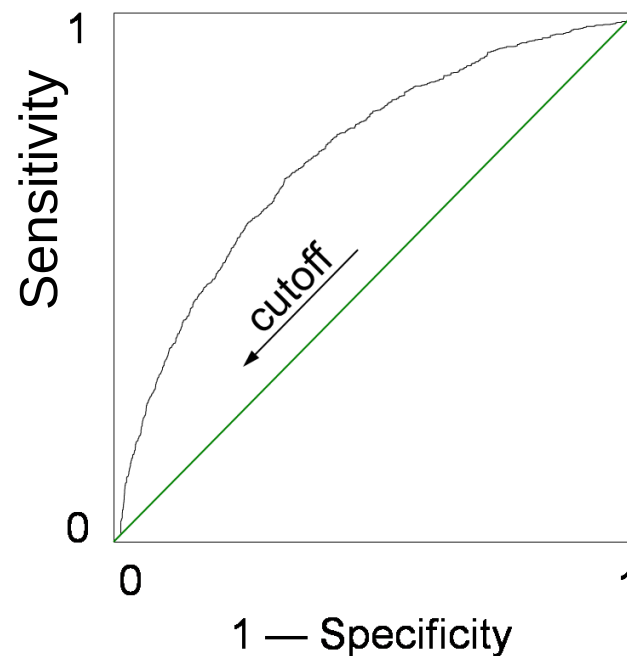
$\beta_1$  = Разница логитов между Low и High

$\beta_2$  = Разница логитов между Medium и High

| Analysis of Maximum Likelihood Estimates |   |    |          |                |                 |            |
|------------------------------------------|---|----|----------|----------------|-----------------|------------|
| Parameter                                |   | DF | Estimate | Standard Error | Wald Chi-Square | Pr > ChiSq |
| Intercept                                |   | 1  | -0.0904  | 0.1608         | 0.3159          | 0.5741     |
| IncLevel                                 | 1 | 1  | -0.6717  | 0.2465         | 7.4242          | 0.0064     |
| IncLevel                                 | 2 | 1  | -0.6659  | 0.2404         | 7.6722          | 0.0056     |

# Оценка модели

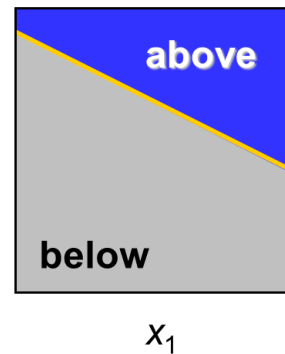
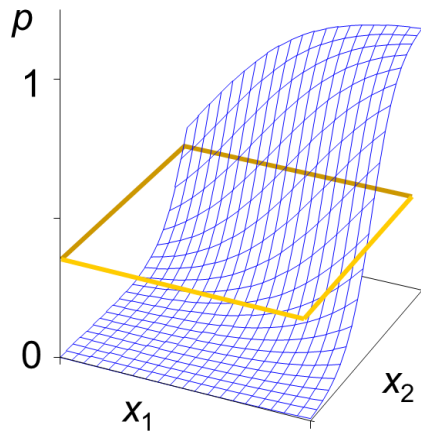
|              |   | Predicted Class       |                       |                 |
|--------------|---|-----------------------|-----------------------|-----------------|
|              |   | 0                     | 1                     |                 |
| Actual Class | 0 | <b>True Negative</b>  | <b>False Positive</b> | Actual Negative |
|              | 1 | <b>False Negative</b> | <b>True Positive</b>  | Actual Positive |
|              |   | Predicted Negative    | Predicted Positive    |                 |



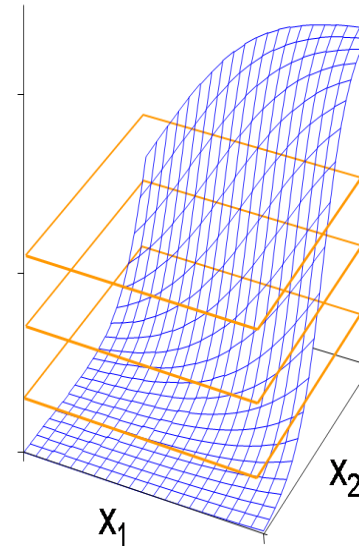
**SENSITIVITY** (true positive rate (TPR),  
hit rate, recall)  
$$TPR = TP / (TP + FN)$$

**SPECIFICITY** (SPC) (true negative  
rate (TNR))  
$$SPC = TN / (FP + TN)$$

# Оценка модели



$x_2$   $P$



|   | 0  | 1  | <u>Se</u> | <u>Sp</u> |
|---|----|----|-----------|-----------|
| 0 | 70 | 5  | .64       | .93       |
| 1 | 9  | 16 |           |           |
| 0 | 66 | 9  | .84       | .88       |
| 1 | 4  | 21 |           |           |
| 0 | 57 | 18 | .96       | .76       |
| 1 | 1  | 24 |           |           |

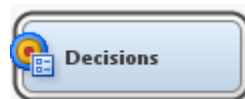
Матрица выигрыша-проигрыша:

|              |   | Decision      |               |
|--------------|---|---------------|---------------|
|              |   | 0             | 1             |
| Actual Class | 0 | $\delta_{TN}$ | $\delta_{FP}$ |
|              | 1 | $\delta_{FN}$ | $\delta_{TP}$ |

Bayes Rule:

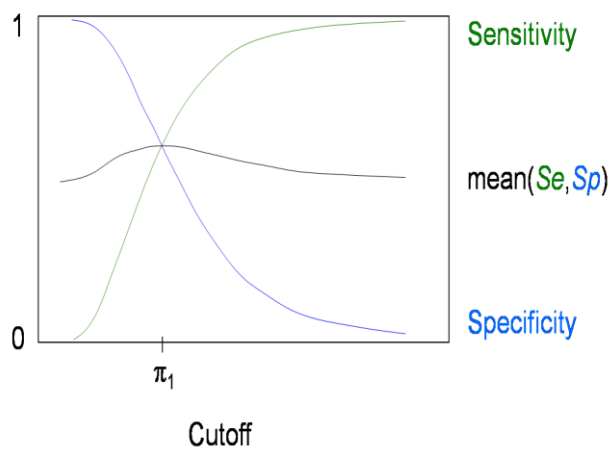
Decision 1 if

$$P > \frac{1}{1 + \left( \frac{\delta_{TP} - \delta_{FN}}{\delta_{TN} - \delta_{FP}} \right)}$$

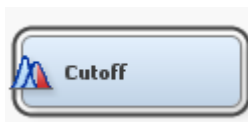
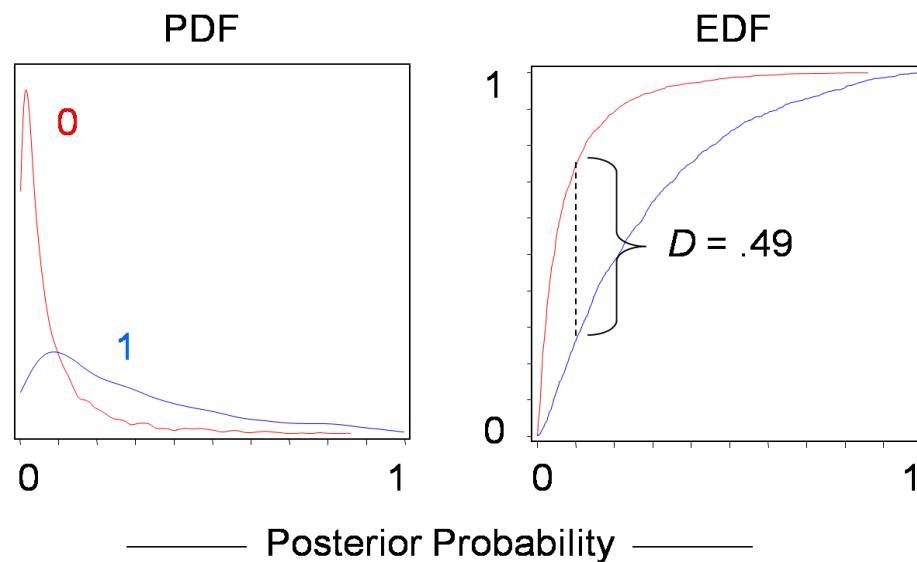


# Выбор порога

Усредненная (иногда взвешенная)  
чувствительность и  
специфичность



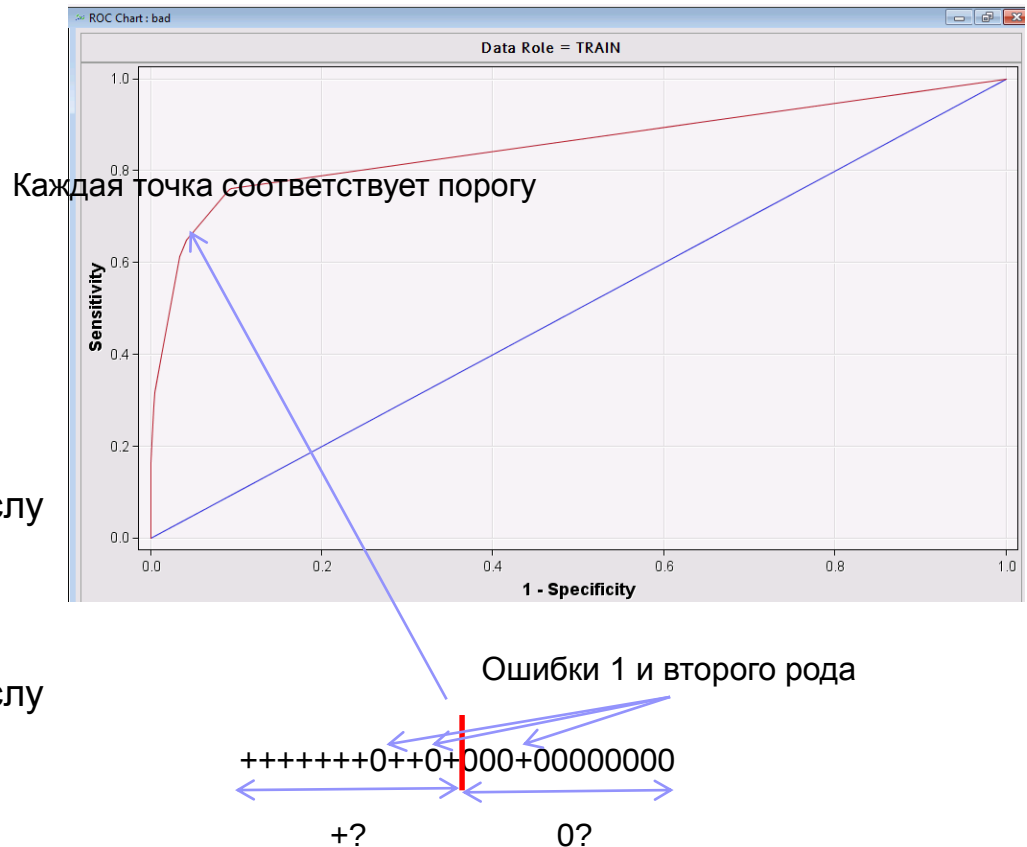
Критерий Колмогорова-Смирнова



# ROC кривая и AUC

## ■ Процедура построения:

- Сортируем набор по убыванию спрогнозированной оценки (вероятности положительного отклика)
- Идем порогом отсечения по отсортированному набору
- Для каждого положения порога считаем:
  1. отношение числа положительных примеров «слева» от порога к числу всех положительных примеров – detection rate
  2. отношение числа отрицательных примеров «слева» от порога к числу всех отрицательных примеров – false positive
- Ставим точку на графике



Оценка на основе согласованности всевозможных пар наблюдений (правильной упорядоченности наблюдений в паре), принадлежащих разным классам.

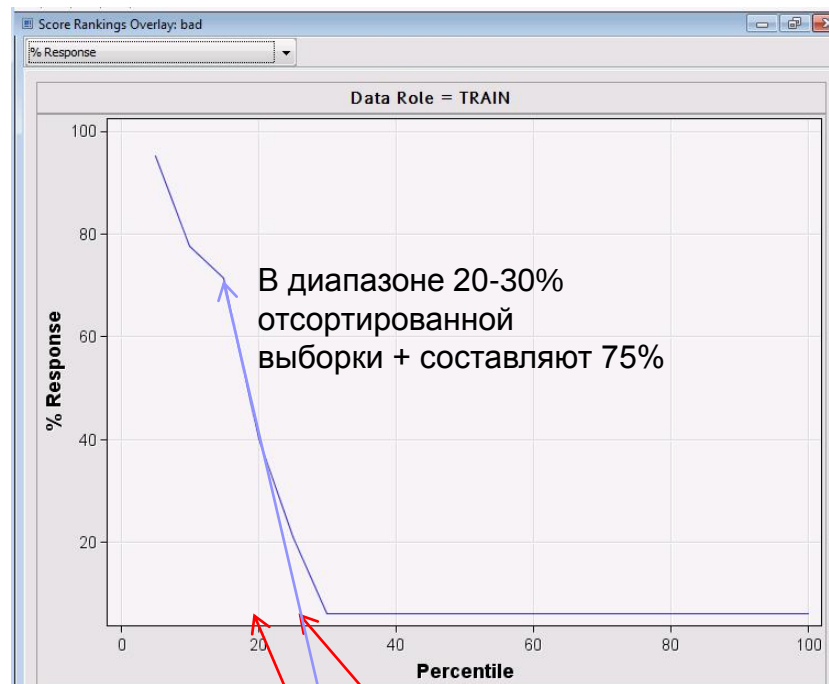
# Графические средства сравнения моделей: Response (отклик)

## ■ Процедура построения:

- Сортируем (например, слева направо) набор по убыванию спрогнозированной оценки (вероятности положительного отклика)
- Идем порогом отсечения по отсортированному набору (слева направо) с некоторым диапазоном (как правило кратно 5%)
- Для каждого положения диапазона считаем отношение числа положительных примеров к числу всех примеров внутри диапазона
- Ставим точку на графике

## ■ Агрегированный отклик (cumulative response):

- Диапазон всегда с 0



++++++0++0+000+00000000

10%

# Графические средства сравнения моделей: Lift (подъем)

## ■ Процедура построения:

- Сортируем (например, слева направо) набор по убыванию спрогнозированной оценки (вероятности положительного отклика)
- Идем порогом отсечения по отсортированному набору (слева направо) с некоторым диапазоном (как правило кратно 5%)
- Для каждого положения диапазона считаем отношение числа положительных примеров к числу положительных примеров, которые могли бы быть выбраны «случайно» - без модели
- Ставим точку на графике

## ■ Агрегированный подъем(cumulative lift):

- Диапазон всегда с 0

