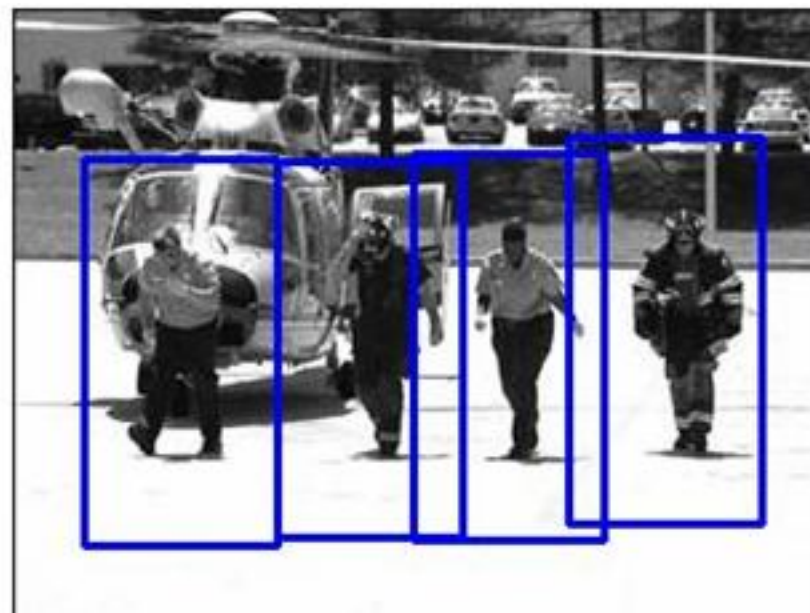
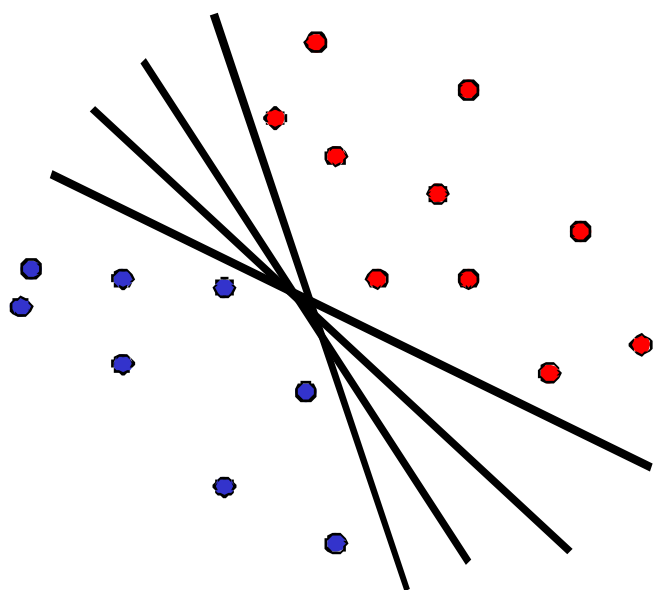




Введение в распознавание изображений, ч. 2



АНТОН КОНУШИН

Slides adapted from Fei-Fei Li, Rob Fergus, Antonio Torralba, Jean Ponce and Svetlana Lazebnik



План

Лекция 3

- Сопоставление шаблонов
- Основы сегментации изображений
- Анализ сегментов

Лекция 4

- Введение в машинное обучение на примере метода опорных векторов
- Алгоритм HOG + SVM для поиска пешеходов на изображении



На прошлой лекции

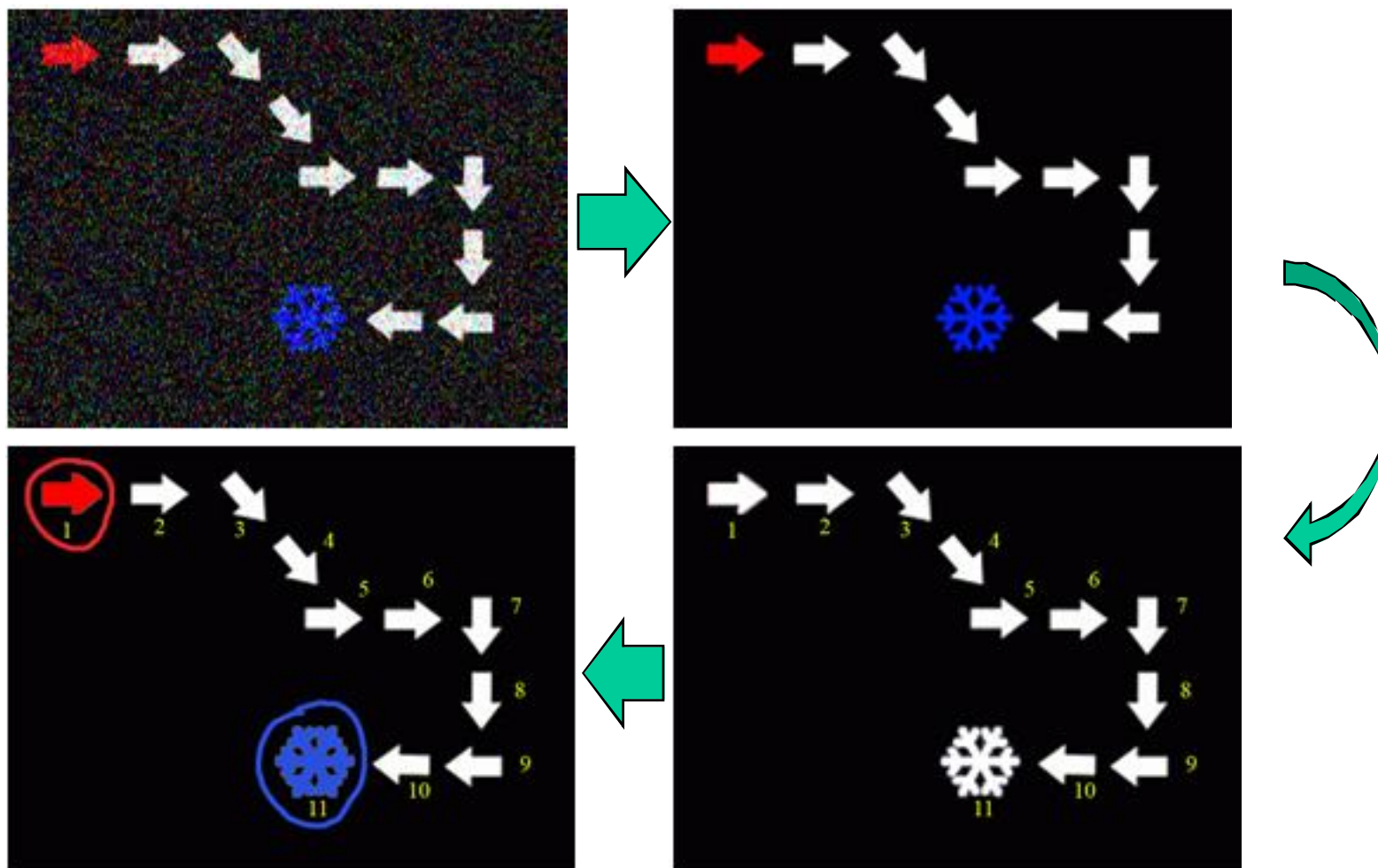
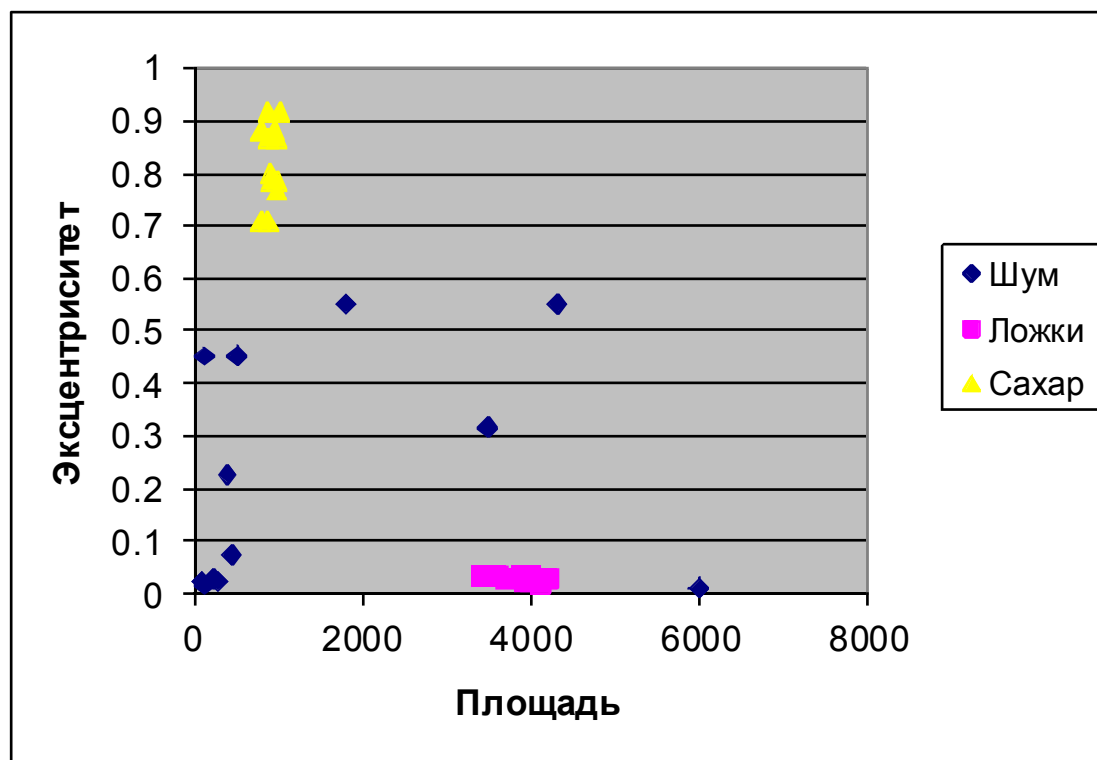


Схема простого алгоритма выделения и распознавания объектов

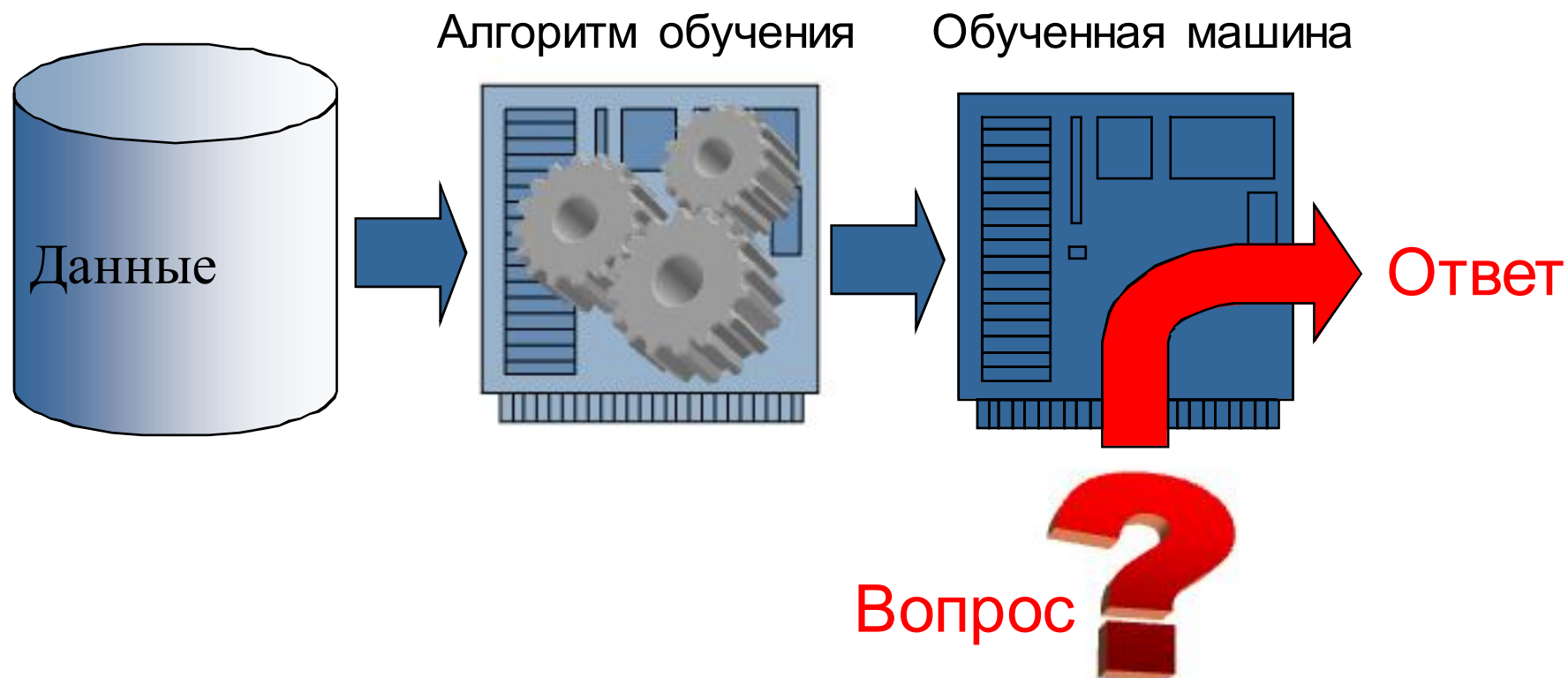


Что было сложно?

- Признаков может быть много (несколько моментов, площадь, положение, гистограммы яркости, каналов цвета и т.д.)
- Подбирать правила приходится вручную, много гипотез, когда признаков много и распределение сложно, это невозможно.



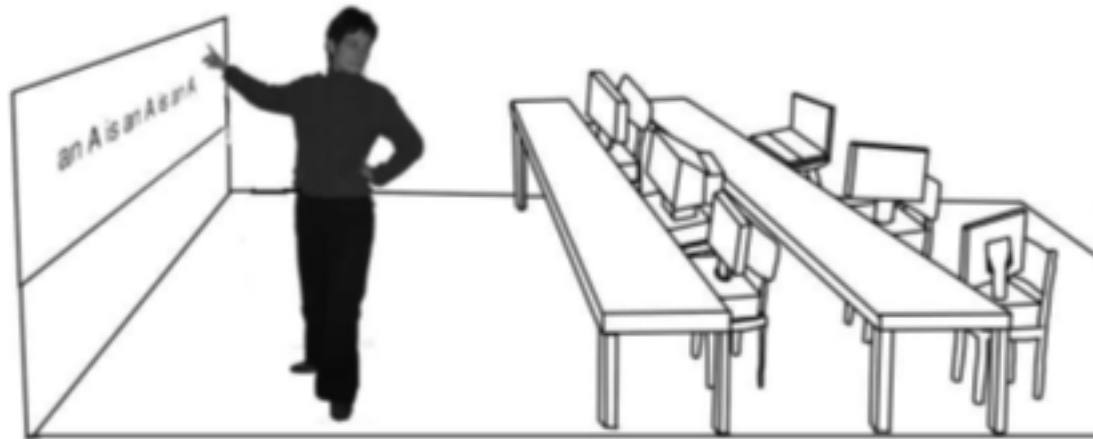
Как должно быть в идеале





Что такое машинное обучение?

- Обучение $\neq$ «заучивание наизусть»
 - Заучить наизусть – для машины не проблема
 - Мы хотим научить машину делать выводы!
 - Машина должна корректно работать на новых данных, которые мы ей раньше не давали
 - По конечному набору обучающих данных машина должна научиться делать выводы на новых данных





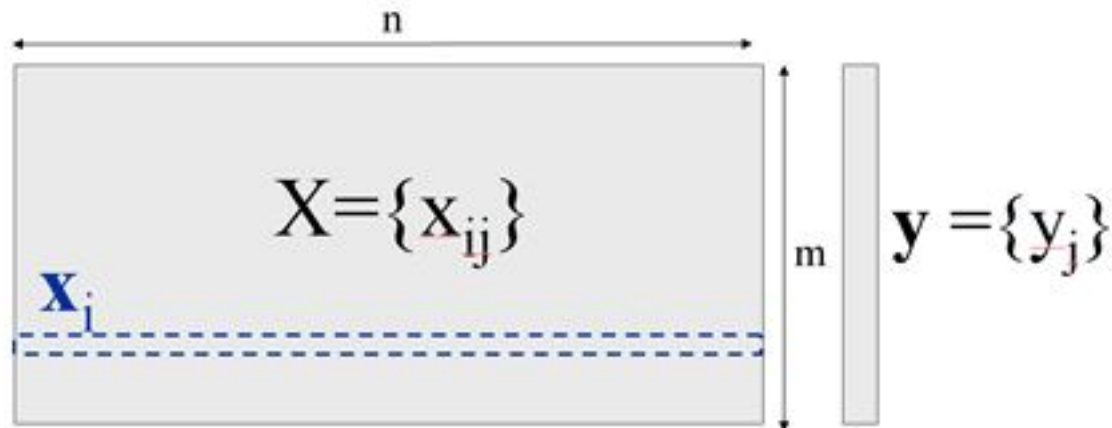
Машинное обучение

- Нейронные сети - Перспептрон (Розенблатт, 1958), когнитрон (Фукушима, 1975)
 - Обратное распространение ошибки (1980е)
 - Метод опорных векторов (1990е)
 - Бустинг (конец 1990х)
 - Рандомизированный решающий лес (начало 2000х)
 - Многослойные нейронные сети и предобучение (2006 и далее)
-
- Мы рассмотрим «на пальцах» один метод – «метод опорных векторов»



Задача классификации образов

- Дано множество объектов, каждый из которых принадлежит одному из нескольких классов. Нужно определить, какому классу принадлежит данный экземпляр
- Каждый объект с номером j можно описать вектором признаков $\mathbf{x}_j$
- Каждому объекту можно приписать метку класса y_j
- Всё множество известных наблюдений (конечное) можно записать в следующем виде:

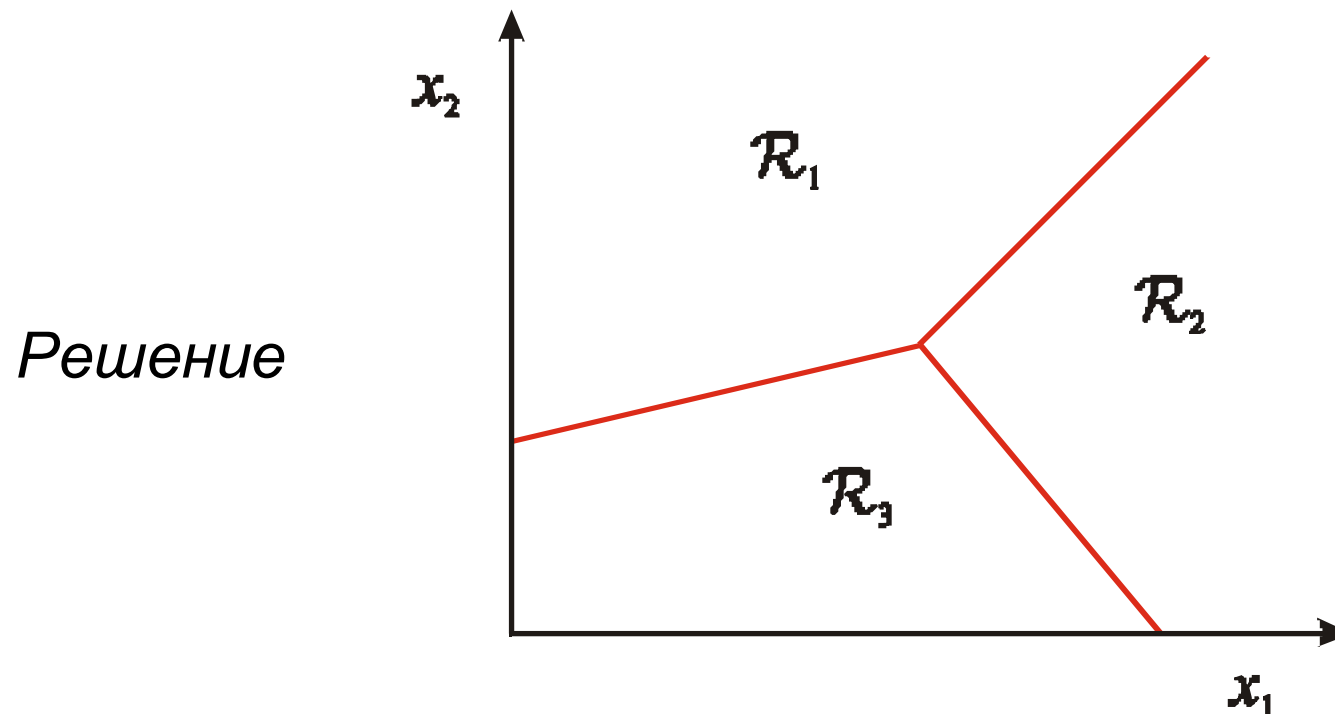


Каждое наблюдение $(\mathbf{x}, y)$ ещё называют прецедентом



Задача классификации образов

- Задача построить функцию $y=f(x)$, которая для каждого вектора признаков x даёт ответ y , какому классу принадлежит объект x
- Функция $f()$ называется *решающее правило* или *классификатор*
- Любое *решающее правило* делит пространство признаков на *решающие регионы* разделенные *решающими границами*





Классификация

- Будем выбирать функции f_t из **параметрического семейства F** (т.е. будем выбирать подходящий набор параметров t)
- Введем некоторую **функцию потерь $L(f_t(x), y)$** ,
 - В случае классификации часто используют $L(f(x), y) = I[y \neq f(x)]$, где $f(x)$ - предсказанный класс
 - Можем использовать и другие функции потерь
- Задача обучения состоит в том, чтобы найти набор параметров классификатора f , при котором потери для *новых* данных будут минимальны
- «Метод классификации» = параметрическое семейство F + алгоритм оценки параметров по обучающей выборке



Общий риск

- Общий риск – математическое ожидание потерь:

$$R(f) = E\left(L\left(f(\mathbf{x}), y\right)\right) = \int_{\mathbf{x}, y} L(f(\mathbf{x}), y) dP$$

- рассчитать невозможно, поскольку распределение P неизвестно



Эмпирический риск

- Дано $X^m = \{\mathbf{x}_1, \dots, \mathbf{x}_m\}$ - обучающая выборка

- Эмпирический риск (ошибка обучения):

$$R_{emp}(f, X^m) = \frac{1}{m} \sum_{i=1}^m L(f(x_i), y_i)$$

- Метод минимизации эмпирического риска:

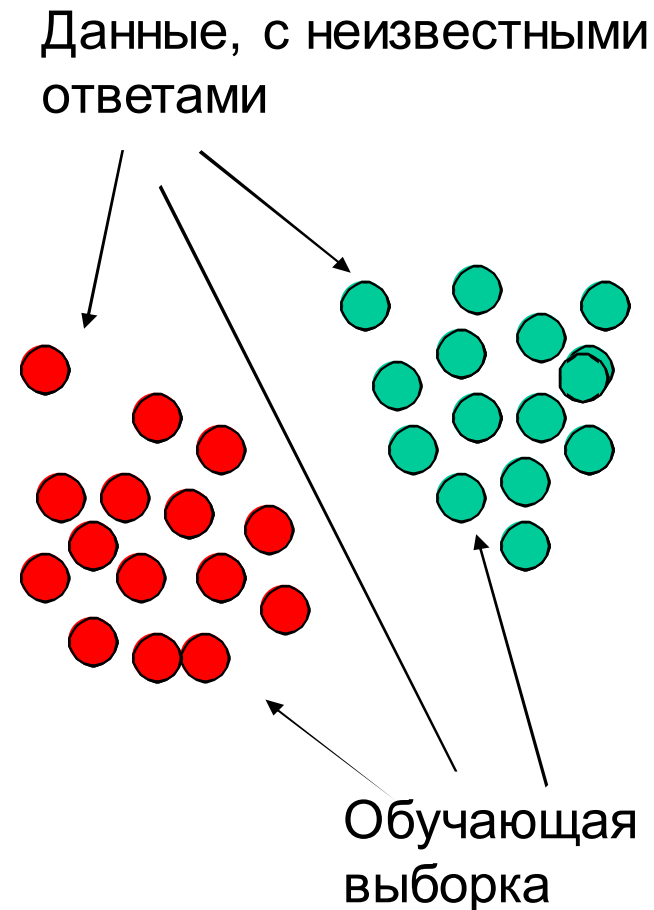
$$f = \arg \min_{f \in F} R_{emp}(f, X^m)$$

Смысл: подбираем параметры $\mathbf{t}$ решающего правила $\mathbf{f}$ таким образом, чтобы эмпирический риск $\mathbf{R}_{emp}$ был минимален



Линейная классификация

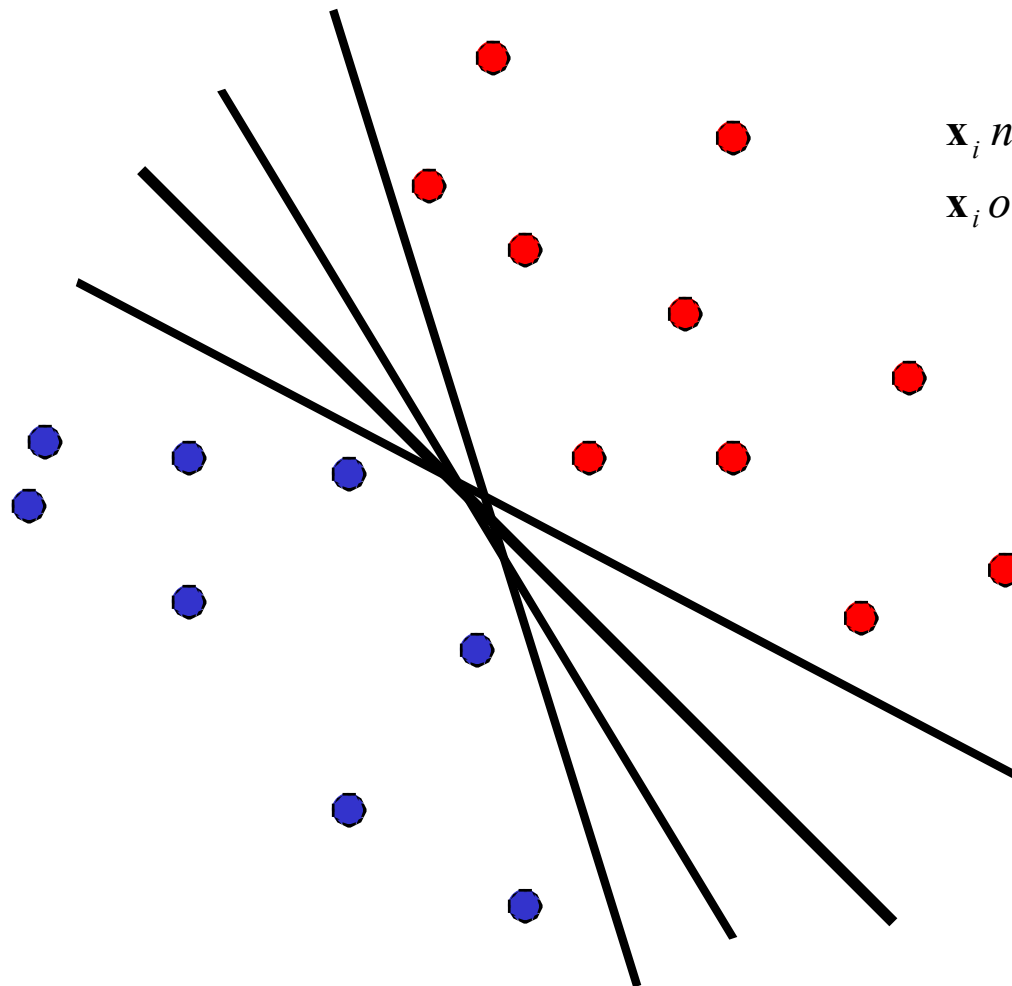
- Рассмотрим случай линейно разделимых данных
- Только для 2х классов
- Т.е. таких, что вектора в пространстве признаков можно отделить друг от друга гиперплоскостью





Линейный классификатор

- Найдем линейную функцию (гиперплоскость) и разделим положительные $\{y=+1\}$ и отрицательные $\{y=-1\}$ примеры



$$x_i \text{ положительные : } x_i \cdot w + b \geq 0$$

$$x_i \text{ отрицательные : } x_i \cdot w + b < 0$$

Для всех
рассмотренных
плоскостей
эмпирический риск
одинаковый

Какая
гиперплоскость
наилучшая?



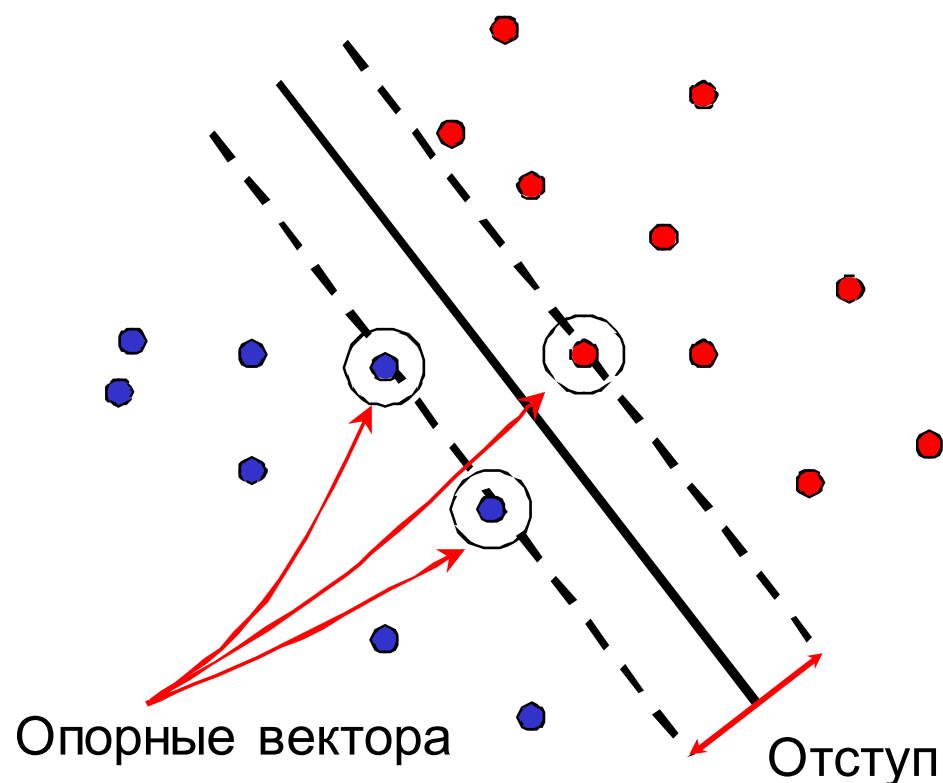
Метод опорных векторов

- Support Vector Machine (SVM)
- Найдем гиперплоскость, максимизирующую *отступ* между положительными и отрицательными примерами



Метод опорных векторов

- Найдем гиперплоскость, максимизирующую *отступ* между положительными и отрицательными примерами



$\mathbf{x}_i$ положительные ($y_i = 1$): $\mathbf{x}_i \cdot \mathbf{w} + b \geq 1$

$\mathbf{x}_i$ отрицательные ($y_i = -1$): $\mathbf{x}_i \cdot \mathbf{w} + b \leq -1$

Для опорных векторов, $\mathbf{x}_i \cdot \mathbf{w} + b = \pm 1$

Расстояние от точки до гиперплоскости: $\frac{|\mathbf{x}_i \cdot \mathbf{w} + b|}{\|\mathbf{w}\|}$

Поэтому отступ равен $2 / \|\mathbf{w}\|$



Поиск гиперплоскости

1. Максимизируем $2/\|\mathbf{w}\|$
2. Правильно классифицируем все данные:

$$\mathbf{x}_i \text{ позитивные } (y_i = 1): \quad \mathbf{x}_i \cdot \mathbf{w} + b \geq 1$$

$$\mathbf{x}_i \text{ негативные } (y_i = -1): \quad \mathbf{x}_i \cdot \mathbf{w} + b \leq -1$$

- Квадратичная оптимизационная задача:

- Минимизируем $\frac{1}{2} \mathbf{w}^T \mathbf{w}$
При условии $y_i(\mathbf{w} \cdot \mathbf{x}_i + b) \geq 1$

- Решается с помощью методом множителей Лагранжа



Поиск гиперплоскости

- Решение: $\mathbf{w} = \sum_i \alpha_i y_i \mathbf{x}_i$

Обученные
веса

Опорные
вектора

- Для большей части векторов вес = 0!
- Все вектора, для которых вес > 0 называются опорными
- Определяется только опорными векторами



Поиск гиперплоскости

- Решение: $\mathbf{w} = \sum_i \alpha_i y_i \mathbf{x}_i$

$b = y_i - \mathbf{w} \cdot \mathbf{x}_i$ для любого опорного вектора

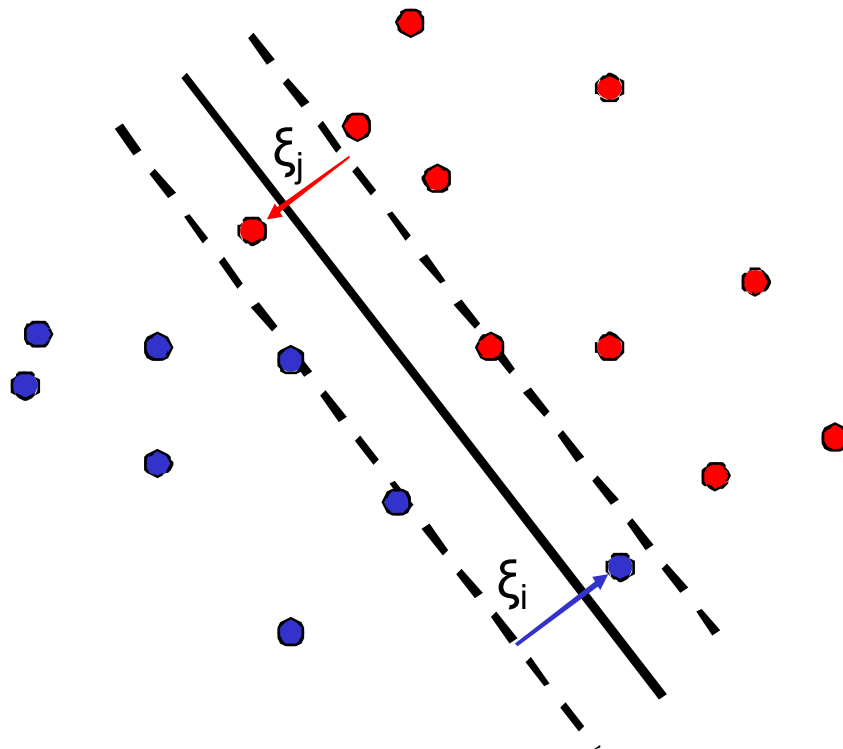
- Решающая функция:

$$\mathbf{w} \cdot \mathbf{x} + b = \sum_i \alpha_i y_i \mathbf{x}_i \cdot \mathbf{x} + b$$

- Решающая функция зависит от скалярных произведений (inner product) от тестового вектора $\mathbf{x}$ и опорных векторов $\mathbf{x}_i$
- Решение оптимизационной задачи также требует вычисления скалярных произведений $\mathbf{x}_i \cdot \mathbf{x}_j$ между всеми парами векторов из обучающей выборки



Реальный случай



Вводим дополнительные «slack»
переменные: $\xi = (\xi_1, \dots, \xi_n)^T$

Минимизируем $\frac{1}{2} \mathbf{w}^T \mathbf{w} + C \sum_{i=1}^n \xi_i$

При условии $y_i (w x_i + b) \geq 1 - \xi_i$

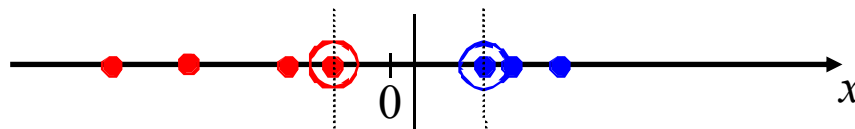
C – параметр регуляризации

В реальном случае идеально разделить данные невозможно
(эмпирический риск всегда больше 0)

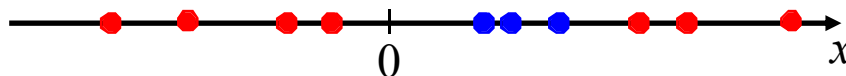


Нелинейные SVM

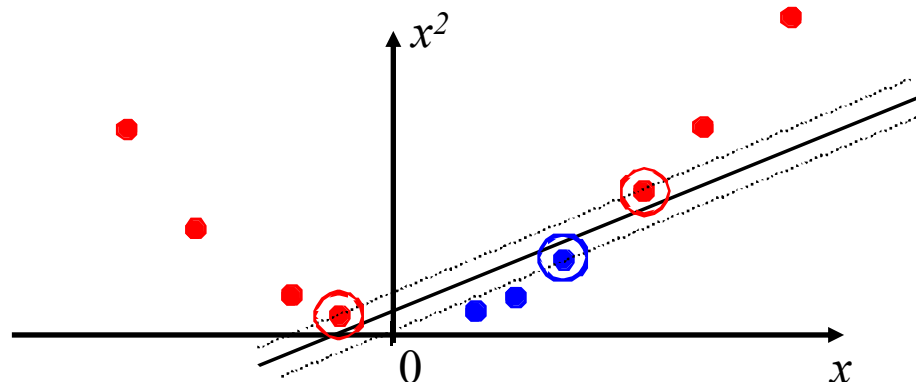
- На линейно разделимых данных SVM работает отлично:



- Но на более сложных данных не очень:



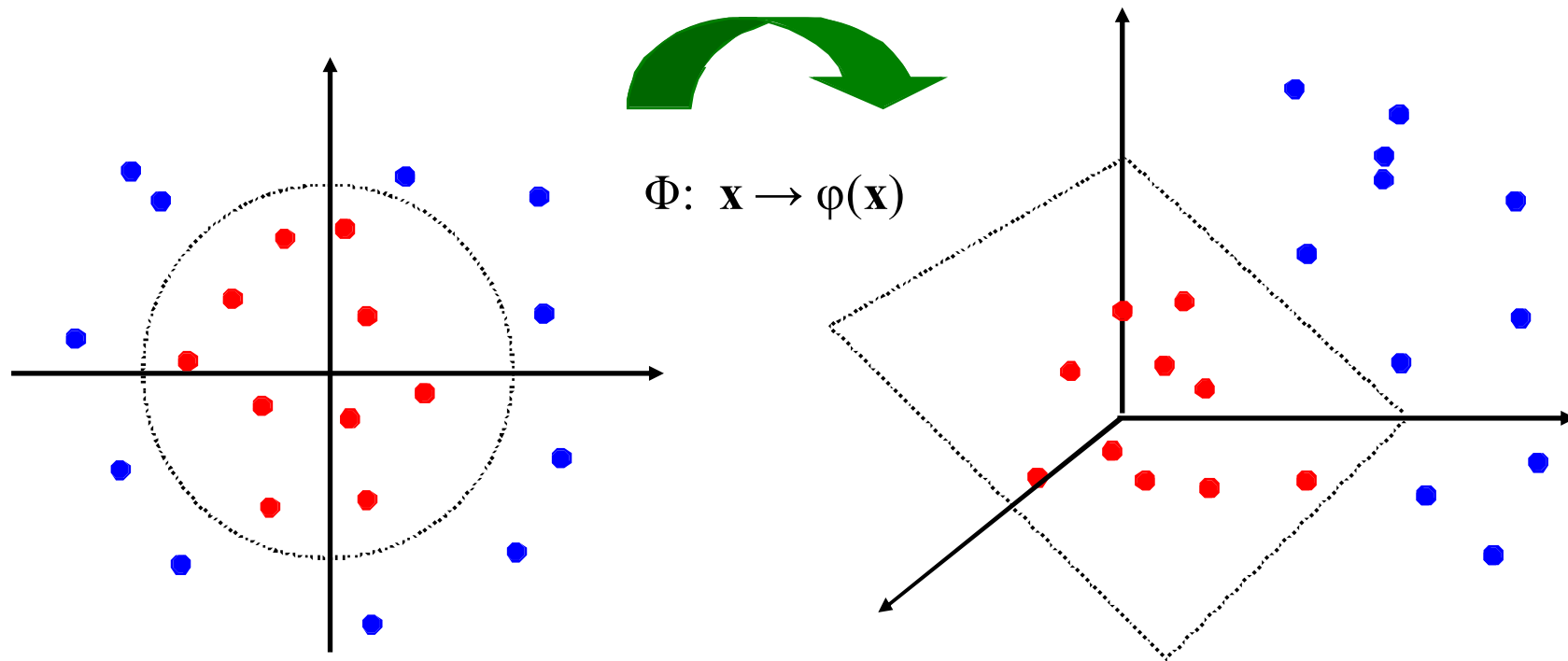
- Можно отобразить данные на пространство большей размерности и разделить их линейно там:





Нелинейные SVM

- **Идея:** отображение исходного пространства параметров на какое-то многомерное пространство признаков (feature space) где обучающая выборка линейно разделима:





Нелинейные SVM

- Вычисление скалярных произведений в многомерном пространстве вычислительно сложно
- *The kernel trick*: вместо прямого вычисления преобразования $\phi(\mathbf{x})$, мы определим ядровую функцию K :

$$K(\mathbf{x}_i, \mathbf{x}_j) = \phi(\mathbf{x}_i) \cdot \phi(\mathbf{x}_j)$$

- Чтобы все было корректно, ядро должно удовлетворять условию Мерсера (*Mercer's condition*)
 - Матрица $K(\mathbf{x}_i, \mathbf{x}_j)$ должна быть неотрицательно определенной
- С помощью ядра мы сможем построить нелинейную решающую функцию в исходном пространстве:

$$\sum_i \alpha_i y_i K(\mathbf{x}_i, \mathbf{x}) + b$$



Пример ядра

- Полиномиальное: $K(x, y) = ((\mathbf{x}, y) + c)^d$
- Пусть $d=2$, $x=(x_1, x_2)$:

$$K(x, y) = ((\mathbf{x}, y) + c)^2 = (x_1 y_1 + x_2 y_2 + c)^2 = \varphi(x) \cdot \varphi(y)$$

$$\varphi(x) = (x_1^2, x_2^2, \sqrt{2}x_1 x_2, \sqrt{2c}x_1, \sqrt{2c}x_2, c)$$

- Т.е. из 2х мерного в 6и мерное пространство



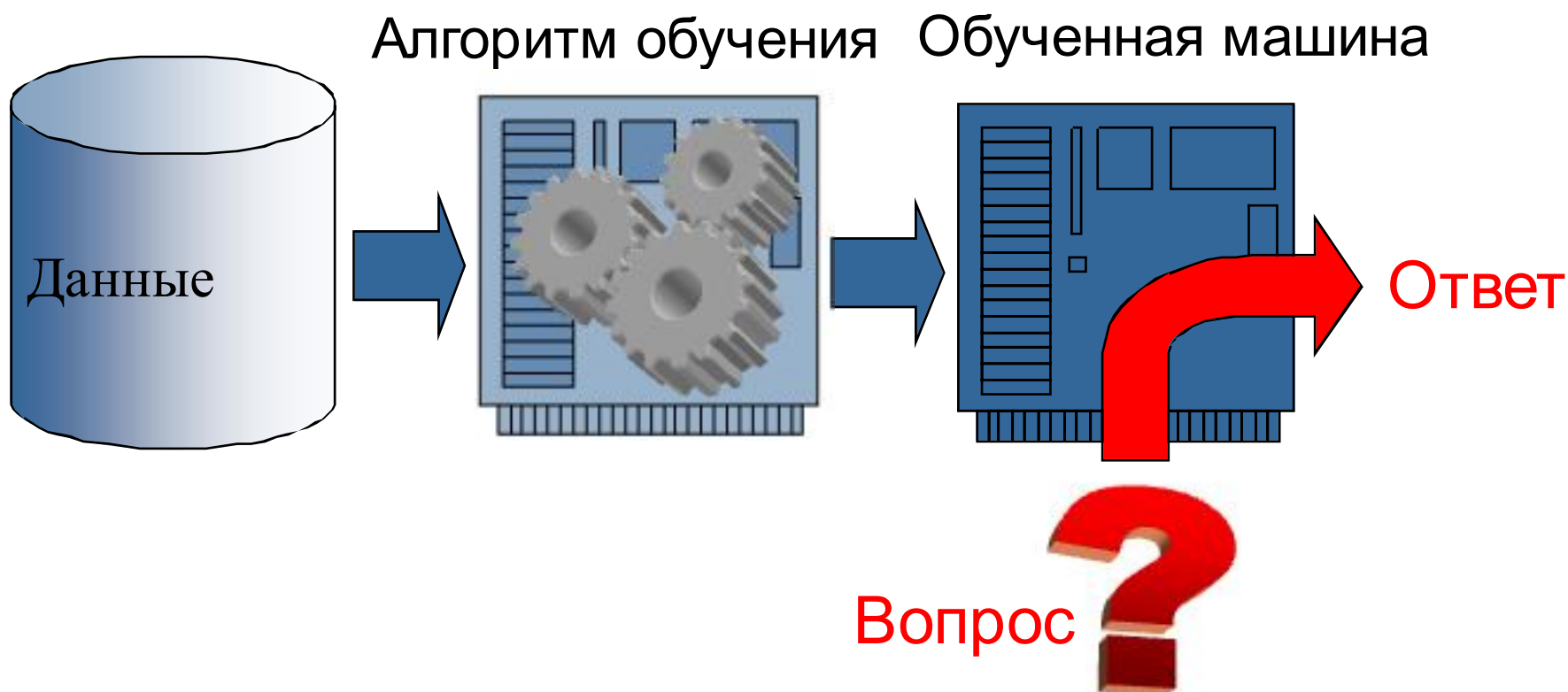
SVM - резюме

- Мощный и гибкий подход на основе ядер
 - Придумано много разных ядер
- На практике работает очень хорошо, даже для маленьких обучающих выборок
- Много доступных библиотек:
<http://www.kernel-machines.org/software>



Оценка классификаторов

- Обучить классификатор SVM умеем
- Как же оценить, насколько хорошо он обучился?
- Нужны количественные оценки качества обучения – предсказательной способности классификатора





Предсказательная способность

- Нам нужно, чтобы классификатор хорошо работал на всех данных
- Нужно минимизировать «общий риск»:

$$R(f, X) = P_{X_m} (f(x) \neq y) = \int_X P(x) [f(x) \neq y] dx$$

- Напрямую его посчитать невозможно, т.к. требует вычислений на неограниченном множестве
- Минимизируем эмпирический риск (ошибку на имеющейся конечной обучающей выборке)
- Как оценить *предсказательную способность*?
 - «Удерживание»
 - «Скользкий контроль» (Cross-validation)



Общая идея

- Будем разбивать имеющуюся обучающую выборку на части
 - На одной части будем обучать классификатор (минимизировать эмпирический риск)
 - На другой будем измерять ошибки обученного классификатора (оценивать предсказательную способность)

$$R(f, X) \sim P(f(x) \neq y \mid X^c) = \frac{1}{c} \sum_{j=1}^c [f(x_j) \neq y_j]$$



Удерживание

- Пусть, имеется набор данных $X^k = \{x_1, \dots, x_k\}$ с известными ответами
- Разобьем $X^l \cup X^c = X^k : X^l \cap X^c = \emptyset$
- Будем использовать для обучения X^l , а для контроля X^c
- То есть: $P(f(x) \neq y) \approx P(f(x) \neq y | X^c)$

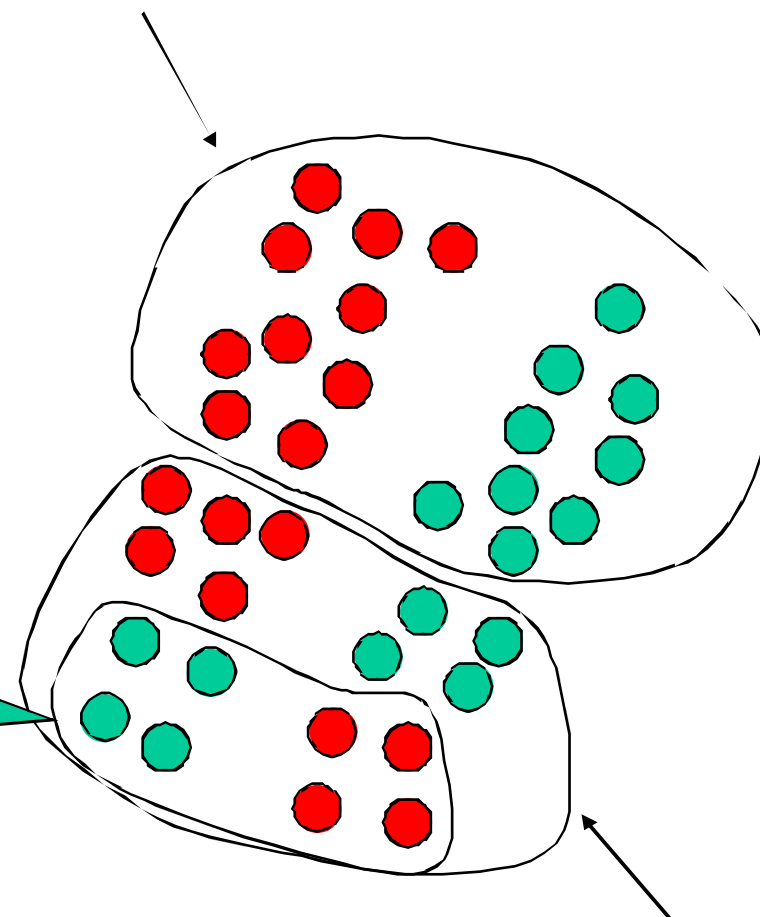


Свойства «удерживания»

- Быстро и просто рассчитывается
- Некоторые «сложные» прецеденты могут полностью попасть в только одну из выборок и тогда оценка ошибки будет смещенной

Ошибка произойдет не по вине классификатора, а из-за разбиения!

Обучение



Контроль



Скольльзящий контроль

- Разделим выборку на d непересекающихся частей и будем поочередно использовать одно из них для контроля а остальные для тренировки

- Разбиваем: $\{X^i\}_1^d : X^i \cap X^j = \emptyset, i \neq j$

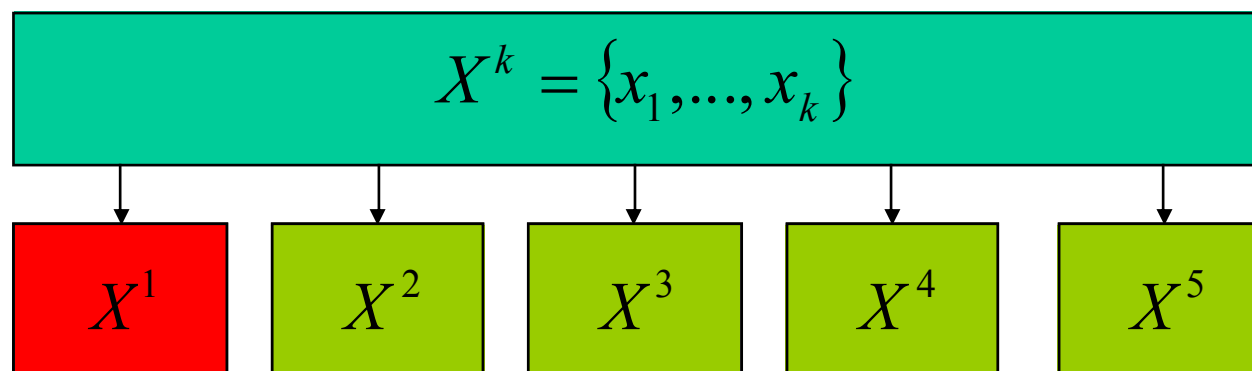
$$\bigcup_{i=1}^d X^i = X^k$$

- Приближаем риск:

$$P(f(X^k) = y^*) \approx \frac{1}{d} \sum_{i=1}^d P(f(X^i) \neq y^* | \bigcup_{i \neq j} X^i)$$



Иллюстрация



Контроль

Результат считается как
средняя



Обучение

ошибка по всем
итерациям



Свойства

- В пределе равен общему риску
- Каждый прецедент будет один раз присутствовать в контрольной выборке
- Обучающие выборки будут сильно перекрываться (чем больше сегментов, тем больше перекрытие)
 - Если одна группа «сложных прецедентов» попала полностью в один сегмент, то оценка будет смещенной
- Предельный вариант – “leave one out”
 - Обучаемся на всех элементах, кроме одного
 - Проверяем на одном
 - Повторяем для всех элементов

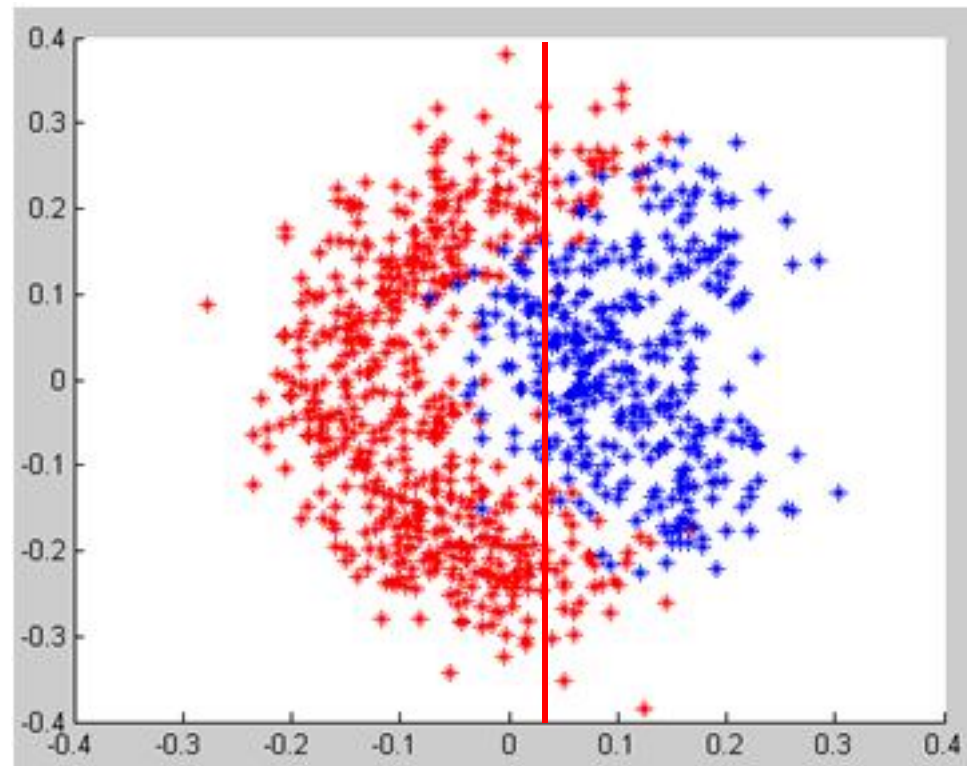
Пример



- Данные – точки на плоскости
- «Классификатор» – порог по оси X

$$a(x^1, x^2) = \begin{cases} +1, x_1 > \Theta \\ -1, x_1 \leq \Theta \end{cases}$$

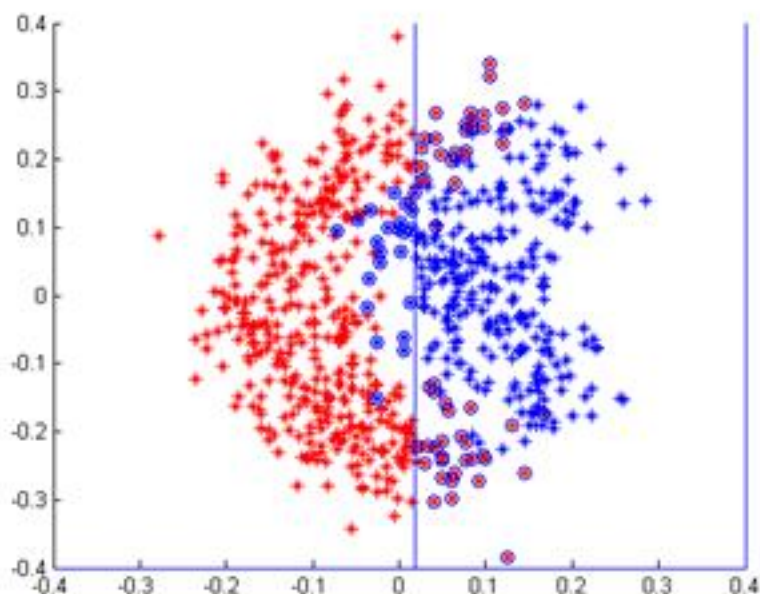
- Будем обучать его, пользуясь разными подходами и измерять ошибки
 - Удерживание
 - Скользящий контроль



Удерживание

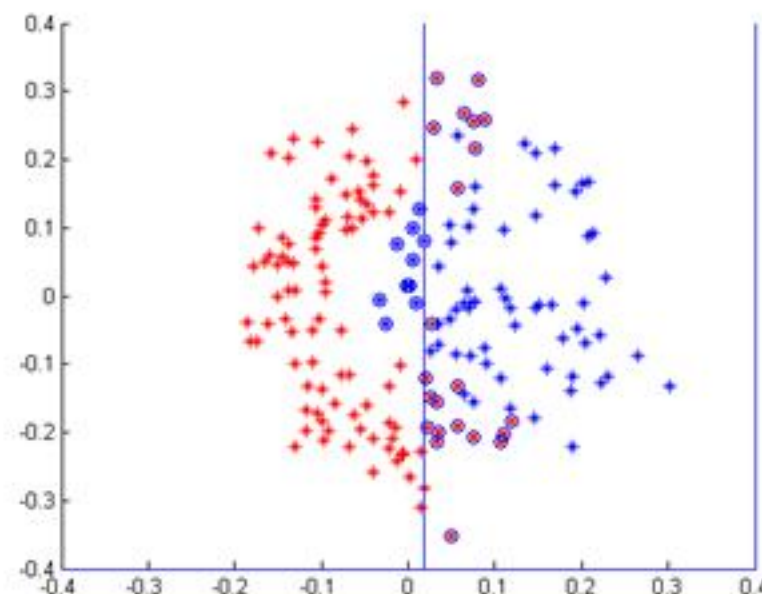


Ошибка: 0.1133



Тренировочная выборка

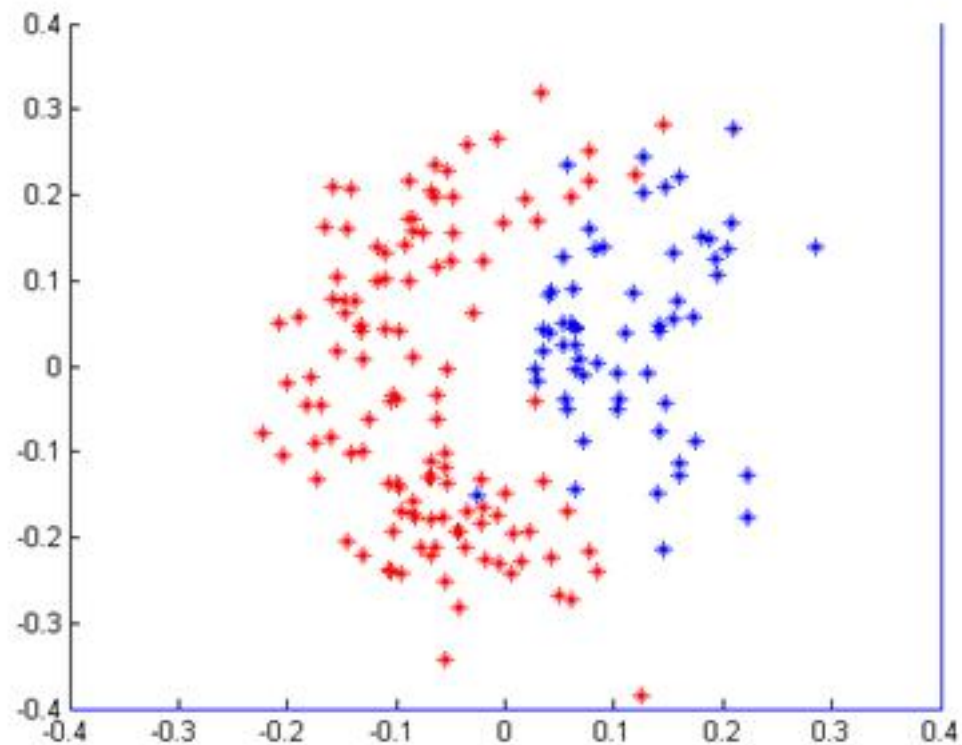
Ошибка: 0.1433



Контрольная выборка

Разобъём данные на 2 части, обучим на одной,
проверим на другой

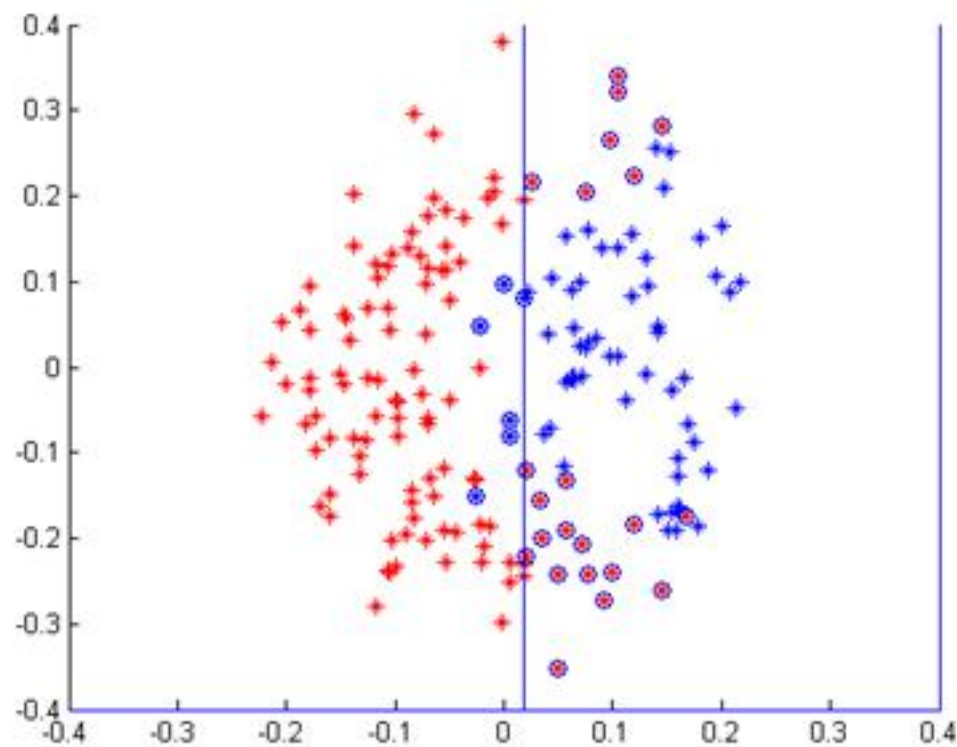
Скользящий контроль: разбиение



Разобъём данные на 5 частей



Скользящий контроль: измерение



Обучим и измерим ошибку 5 раз



Скользющий контроль

- Тренировочная ошибка:
 - $\{0.1236 \quad 0.1208 \quad 0.1250 \quad 0.1097 \quad 0.1306\}$
 - Среднее = 0.1219
- Ошибка на контроле
 - $\{0.1500 \quad 0.1333 \quad 0.1222 \quad 0.1778 \quad 0.1000\}$
 - Среднее = 0.1367



Виды ошибок

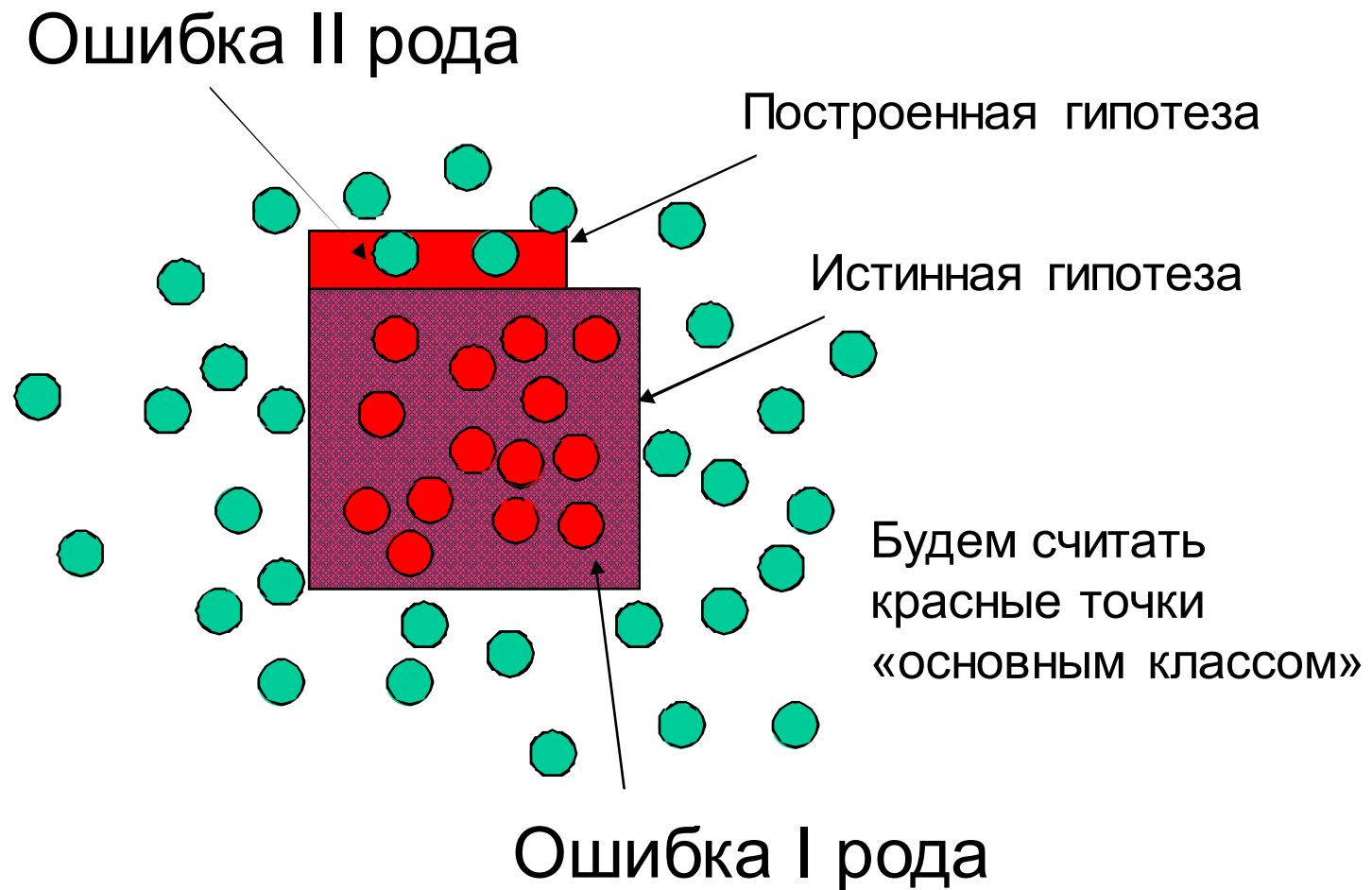
- Измерения ошибки как «вероятности выдать неверный ответ» может быть не всегда достаточно
 - 15% ошибки при постановке диагноза может означать как и то что, 15 % больных будут признаны здоровыми (и возможно умрут от отсутствия лечения), так и то, что 15% здоровых больными (и деньги на лечение будут потрачены зря)
- При неравнозначности ошибок для разных классов вводят понятие ошибки первого и второго рода и замеряют их по отдельности



Ошибки I и II рода

- Пусть, существует «основной класс»
 - Обычно, это класс, при обнаружении которого, предпринимается какое-либо действие;
 - Например, при постановке диагноза основным классом будет «болен», а вторичным классом «здоров».
- **Ошибка первого рода** равна вероятности принять основной класс за вторичный
 - Вероятность «промаха», когда искомый объект будет пропущен
- **Ошибка второго рода** равна вероятности принять вторичный класс за основной
 - Вероятность «ложной тревоги», когда за искомый объект будет принят «фон»

Ошибки I и II рода





Ошибки I и II рода

- Что считать основным классом зависит полностью от специфики прикладной задачи
 - В компьютерном зрении – почти всегда сильно несбалансированы
- Особенно важно оценивать ошибки I и II рода отдельно при несбалансированности классов:
 - Пусть $P(y = +1) = 0.01; P(y = -1) = 0.99$
 - Тогда при ошибке II рода 0 и ошибке I рода 0.5
$$P(f(x) = -1 \mid y = +1) = 0.5$$
 - Общая ошибка всего лишь
$$P(a(x) \neq y) = 0.005$$



Чувствительность vs Избирательность

- **Чувствительность** – вероятность дать правильный ответ на пример основного класса

$$sensitivity = P(f(x) = y \mid y = +1)$$

- Также уровень обнаружения (*detection rate*)

- **Избирательность** – вероятность дать правильный ответ на пример вторичного класса

$$specificity = P(f(x) = y \mid y = -1)$$

- Обычно считают долю ложных обнаружений

$$False\ positive\ rate = 1 - specificity$$



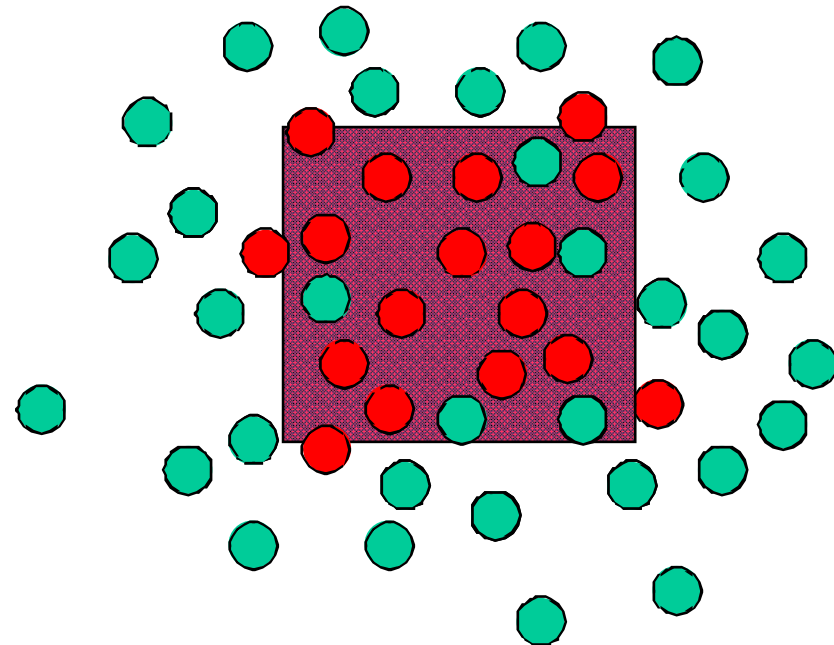
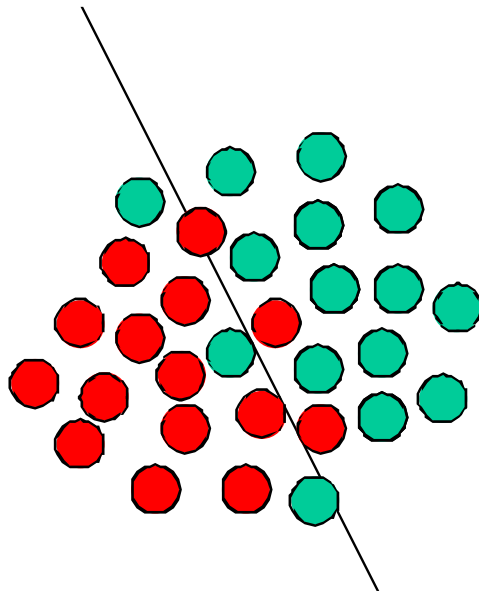
Другие метрики

- Точность (Precision)
 - Доля истинных объектов основного класса среди всех классифицированных, как первый класс
- Полнота (Recall)
 - Доля правильно распознанных объектов основного класса среди всех объектов основного класса из тестовой выборки
- Интегральный показатель F (F-measure)
 - $$F = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$



Регулировка баланса

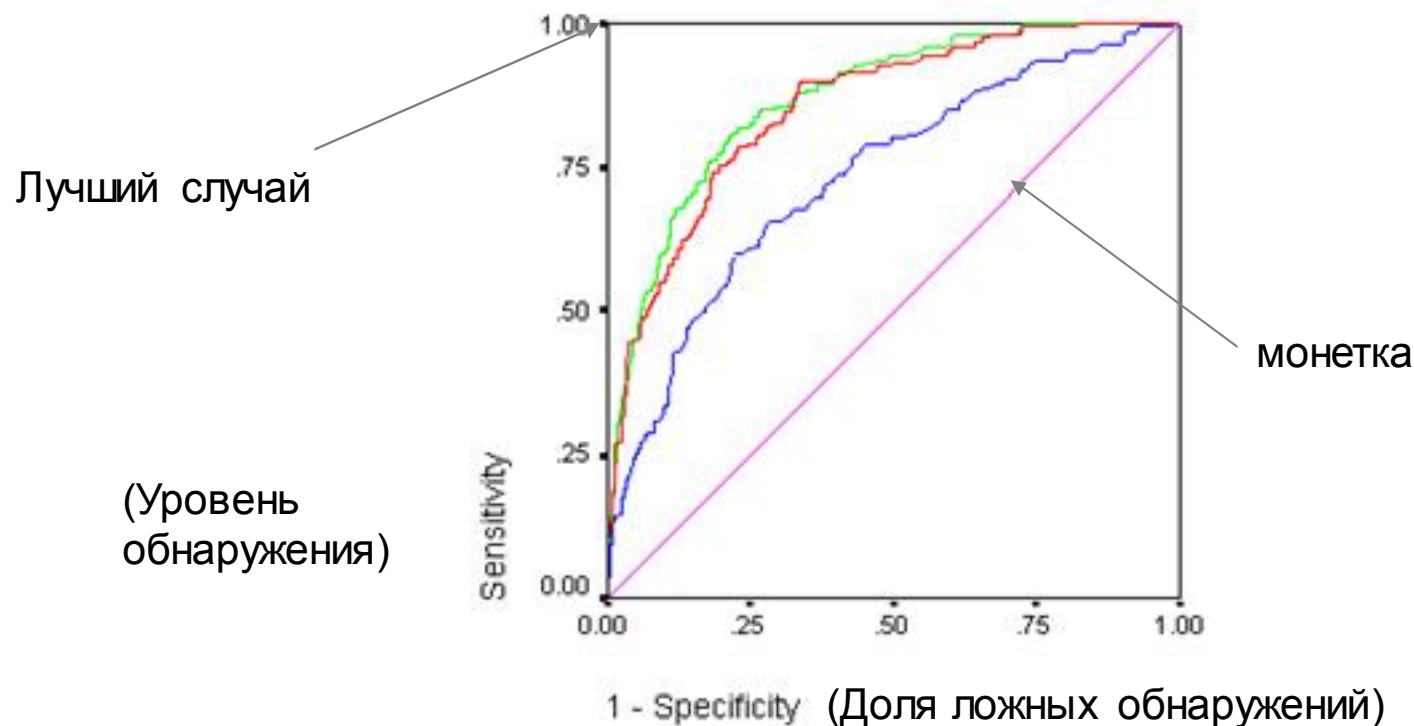
- Почти все алгоритмы классификации допускают регулировку соотношения ошибки I и II рода за счет варьирования некоторого параметра





ROC кривая

- ROC – Receiver Operating Characteristic curve
 - Кривая, отражающая зависимость чувствительности и ошибки второго рода

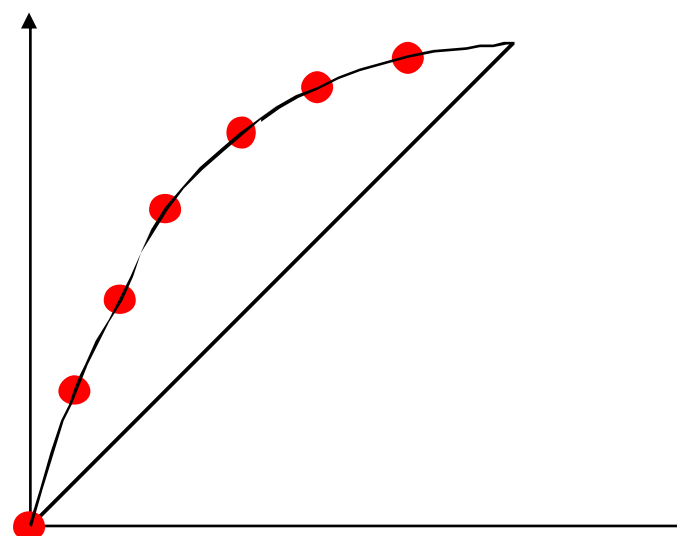




ROC кривая - Построение

- Для различных значений параметра строится таблица ошибок
 - Сам параметр в таблице не участвует!
 - Классификатор строится и оценивается на разных выборках!
- По таблице строится набор точек в плоскости sensitivity x (false positive)
 - Каждая строка таблицы - точка
- По точкам строится кривая

Sensitivity	False Positive
0.0	0.0
0.25	0.5
0.5	0.8
...	...
1.0	1.0

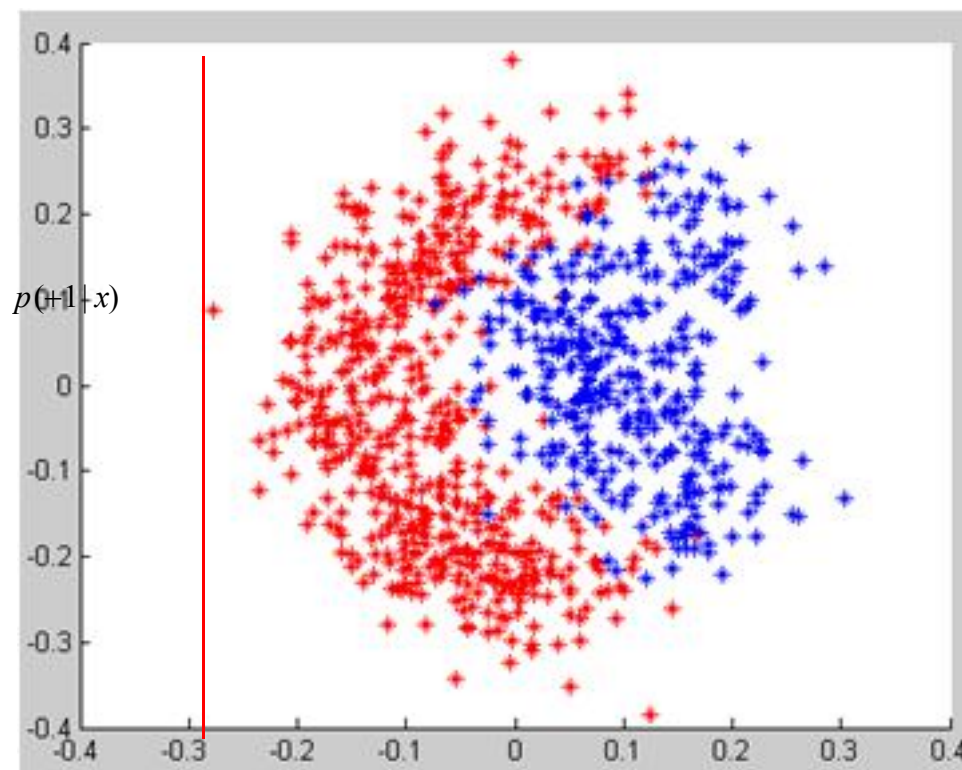




Построение ROC: таблица

- Меняем порог и оцениваем ошибку

Sensitivity	False Positive
0.0	0.0
0.25	0.5
0.5	0.8
...	...
1.0	1.0

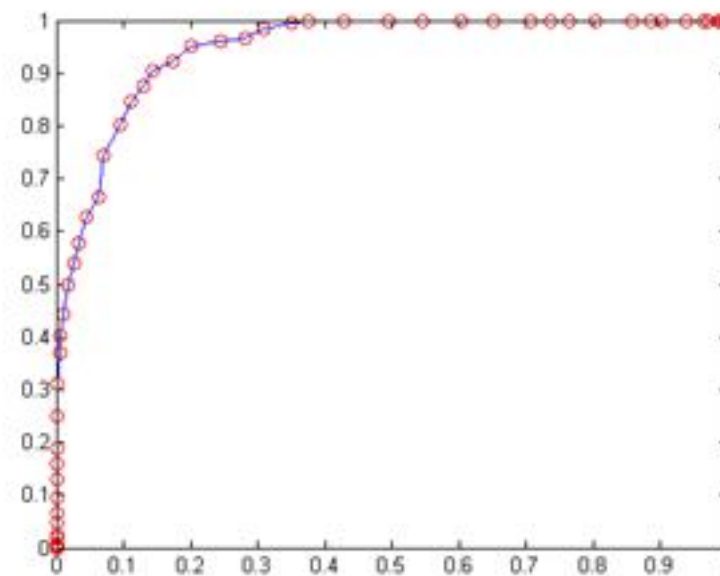


- Теперь построим ROC-кривую
- Поскольку у нас только 1 параметр, будем его менять напрямую
 - Это не обучение, должно быть обучение, но для примера сойдёт



Построение ROC-кривой

- По таблице строим точки
- Точки интерполируем кривой





Анализ ROC кривой

- Площадь под графиком – AUC
 - Дает некоторый объективный показатель качества классификатора
 - Позволяет сравнивать разные кривые
- Соблюдение требуемого значения ошибок I и II рода
 - Зачастую, для конкретной задачи существуют рамки на ошибку определенного рода. С помощью ROC можно анализировать возможность текущего решения соответствовать требованию



Резюме бинарной классификации

- Соберём обучающую выборку
- Подберём оптимальные параметры SVM (ядро, параметры ядра, количество опорных векторов)
 - Зафиксируем параметры
 - Воспользуемся кросс-валидацией, например, с разбиением выборки на 5 блоков
 - 5 раз проведем операцию
 - Выберем блок для контроля
 - Обучим на оставшихся 4х
 - Усредним характеристики (ошибку на контроле, ошибки 1 и 2 рода)
 - Построим ROC-кривую
 - Выберем точку с оптимальными для нашей задачи ошибками



План

Лекция 3

- Сопоставление шаблонов
- Основы сегментации изображений
- Анализ сегментов

Лекция 4

- Введение в машинное обучение на примере метода опорных векторов
- Алгоритм HOG + SVM для поиска пешеходов на изображении



Поиск пешеходов



- Нужно найти «пешеходов» (если есть) и выделить их прямоугольной рамкой
- Пешеход (pedestrian) – это «категория», а не «объект»
 - Люди все похожи, но все разные
 - По простым признакам (цвету, площади) выделить пешехода нельзя



Классификация изображения

- Простая бинарная классификация может ответить на вопрос «да» / «нет»
- Например: «есть» пешеход на изображении или «нет» пешехода на изображении



Обучающая выборка



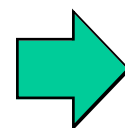
«Картинка ли это пешехода?»

- А нам нужно знать, где ещё на картинке пешеход!



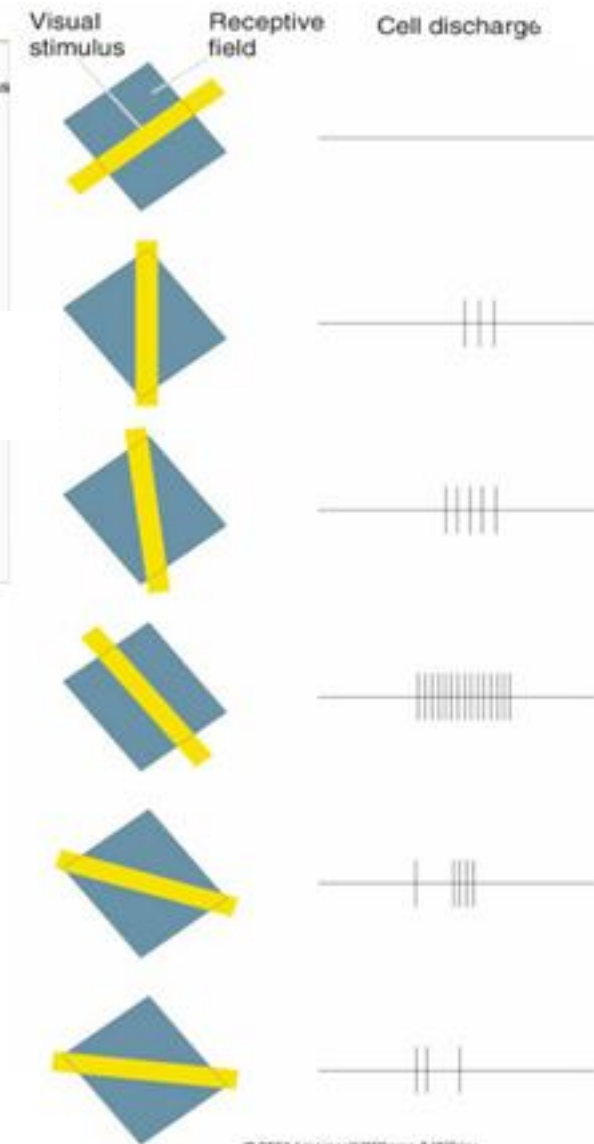
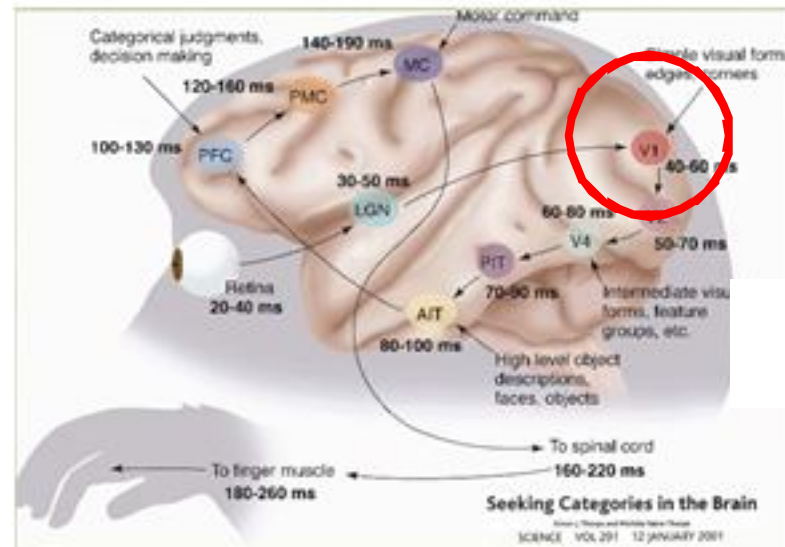
Скользящее окно (*sliding window*)

- Возьмём «окно», пусть размером 64×128 пикселей
- Сканируем изображение «окном»
- Применяем классификатор «пешеход?» к каждому окну
- Скользящее окно – форма сегментации изображения!





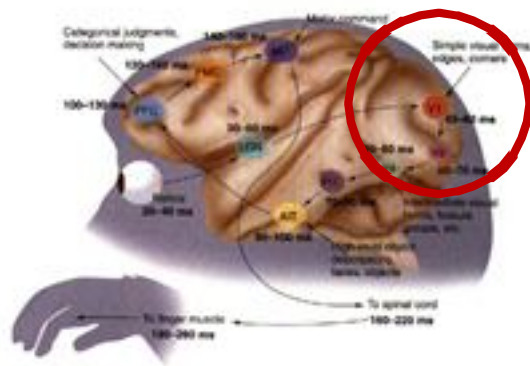
«Простые клетки» V1



Вспомним, что первичная зрительная кора в основном занимается анализом краёв и градиентов в изображении

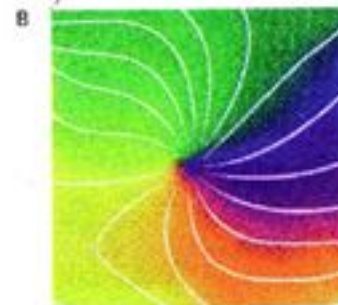
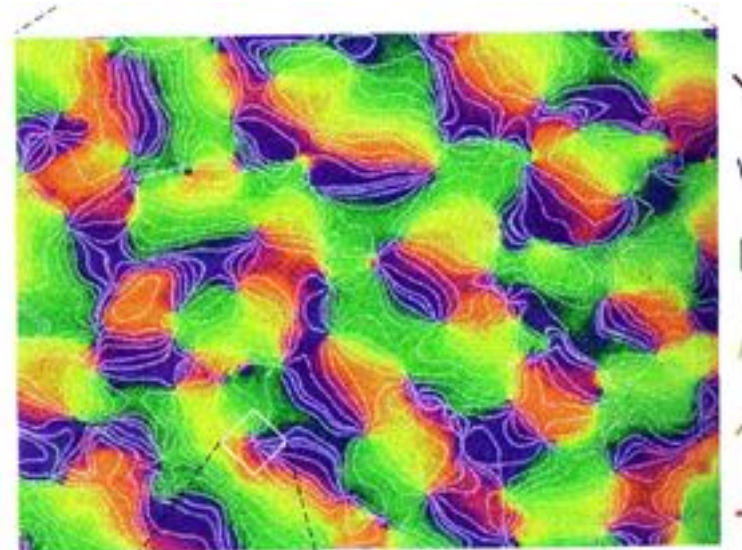


«Ретинотопическая» организация



Для каждой области визуального поля есть клетки, чувствительные к краям разной ориентации (отображены на рисунке цветом)

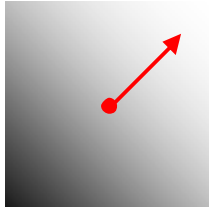
Построим признаки, организованные по похожей схеме!





Оценка градиентов

- Градиент изображения – направление максимального изменения яркости изображения


$$\nabla f = \left[\frac{\partial f}{\partial x}, \frac{\partial f}{\partial y} \right]$$

Сила края:

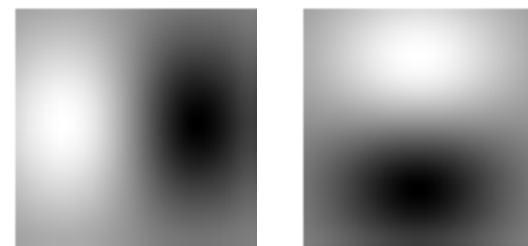
$$\|\nabla f\| = \sqrt{\left(\frac{\partial f}{\partial x}\right)^2 + \left(\frac{\partial f}{\partial y}\right)^2}$$

Направление градиента:

$$\theta = \tan^{-1} \left(\frac{\partial f}{\partial y} / \frac{\partial f}{\partial x} \right)$$

$$\begin{bmatrix} -1 & -2 & -1 \\ 0 & 0 & 0 \\ 1 & 2 & 1 \end{bmatrix} \quad \begin{bmatrix} -1 & 0 & 1 \\ -2 & 0 & 2 \\ -1 & 0 & 1 \end{bmatrix} \quad \text{Собеля}$$

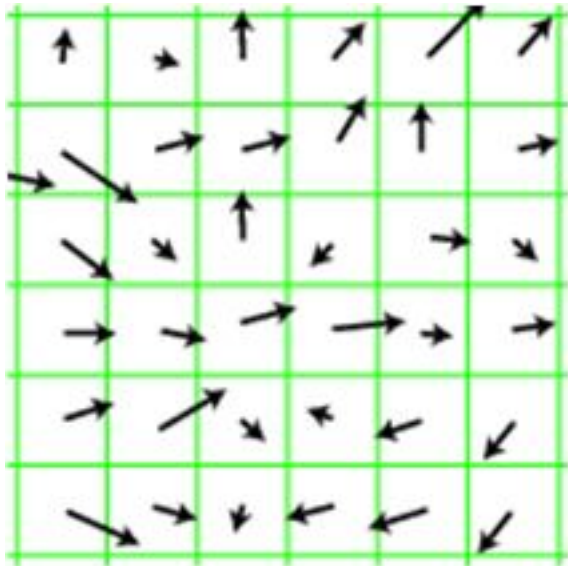
Вычисление разностной
производной через свёртку



Вычисление градиента со
сглаживанием



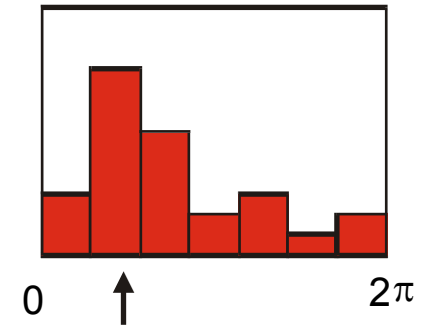
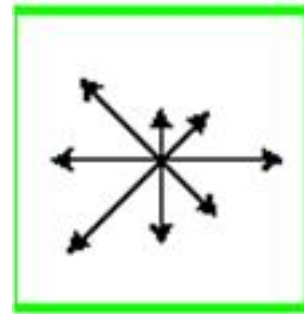
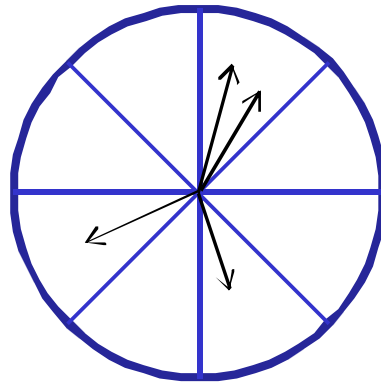
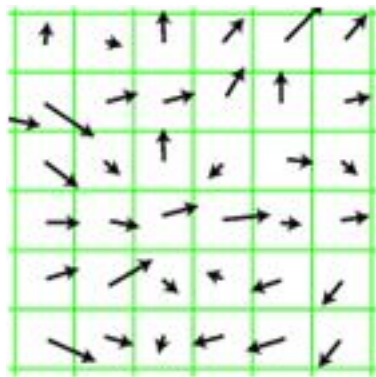
Признаки через градиенты



- Вычислим градиенты в каждом пикселе
 - Направление градиента
 - Силу градиента
- Какие признаки мы можем получить из градиентов?
 - Использовать напрямую (слишком много, неудобно)
 - Посчитать какие-нибудь статистики



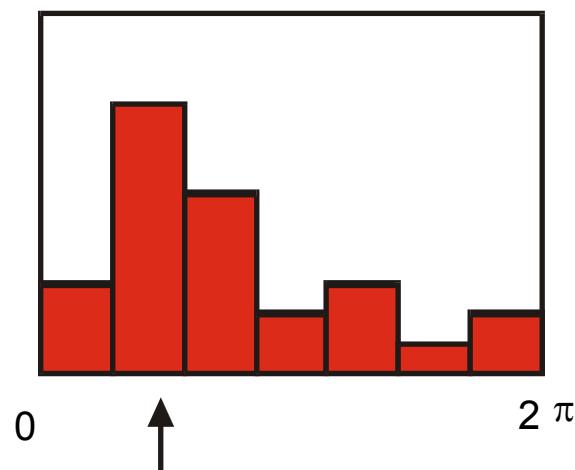
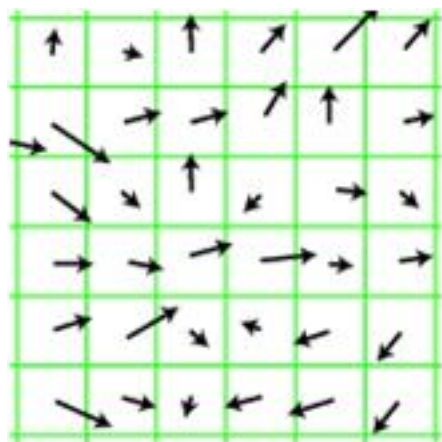
Гистограмма ориентаций градиентов



- Вычислим направление градиента в каждом пикселе
- Квантуем ориентации градиентов на 8 ячеек (направлений)
 - Понетим каждый пиксель номером ячейки
- Посчитаем гистограмму направлений градиентов
 - Для каждой ячейки посчитаем количество пикселей с номером этой ячейки
- Можем нормализовать гистограмму



Гистограмма ориентаций градиентов



- Histogram of oriented gradients (HOG)
- Мощный класс признаков
- Очень широко используется при анализе изображений
- Устойчив к изменениям освещенности изображения
 - Т.к. считаем только направления градиентов

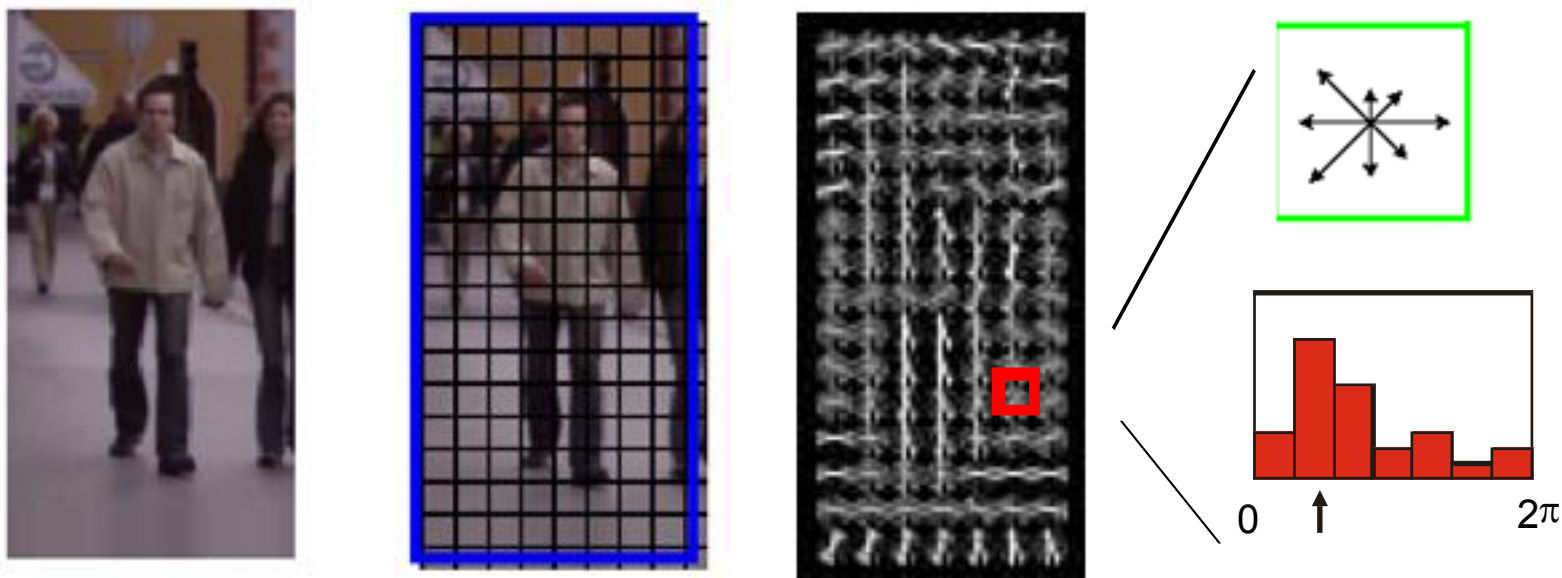


Алгоритм поиска пешеходов

- Алгоритм поиска:
 - Скользящее окно поиска
 - Признаки «гистограммы ориентаций градиентов»
 - Бинарный классификатор «пешеход?» SVM
- Хотя схема HOG + SVM изначально была предложена для пешеходов, она успешно применялась в дальнейшем к разным категориями объектов



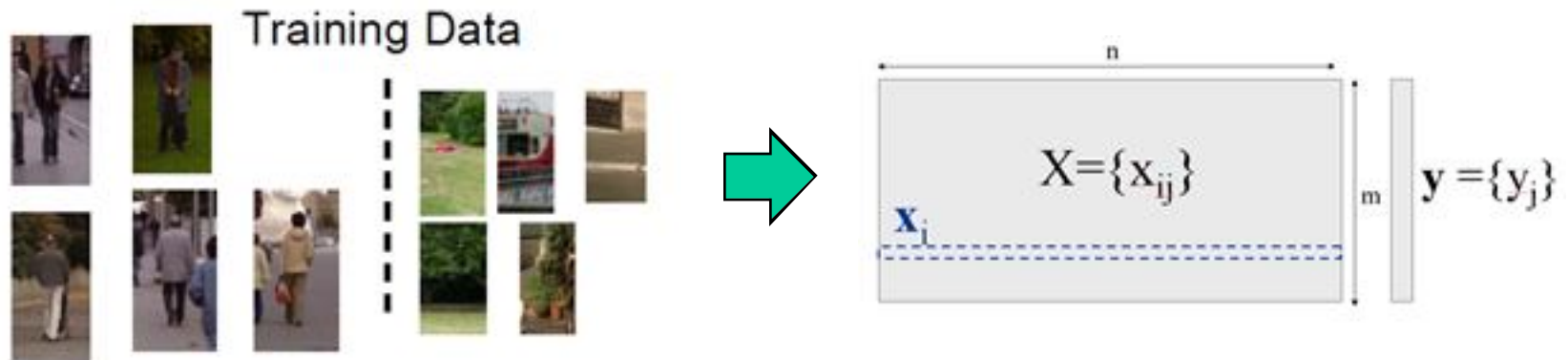
Вычисление признаков



- Возьмём окно 64 x 128 пикселей
- Разобъём его на блоки 8 x 8 пикселей
- Всего будет $8 * 16 = 128$ блоков
- В каждом блоке посчитаем гистограмму ориентаций градиентов с 8 ячейками (8 параметров)
- Всего у нас получится $128 * 8 = 1024$ признака



Обучение классификатора



- Соберём обучающую выборку фрагментом изображения с пешеходами и без
- Для каждого фрагмента посчитаем вектор признаков x и метку y
- На полученной обучающей выборке обучим линейный классификатор SVM



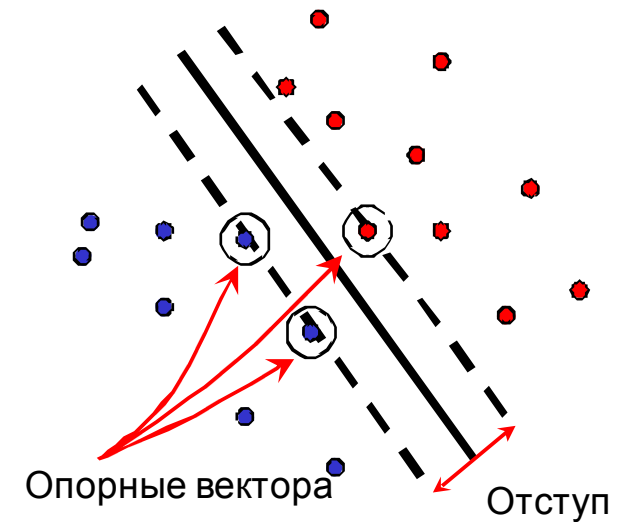
SVM

- Обучаем линейную SVM:

$$f(x) = w^T x + b$$

$$\mathbf{w} = \sum_i \alpha_i y_i \mathbf{x}_i$$

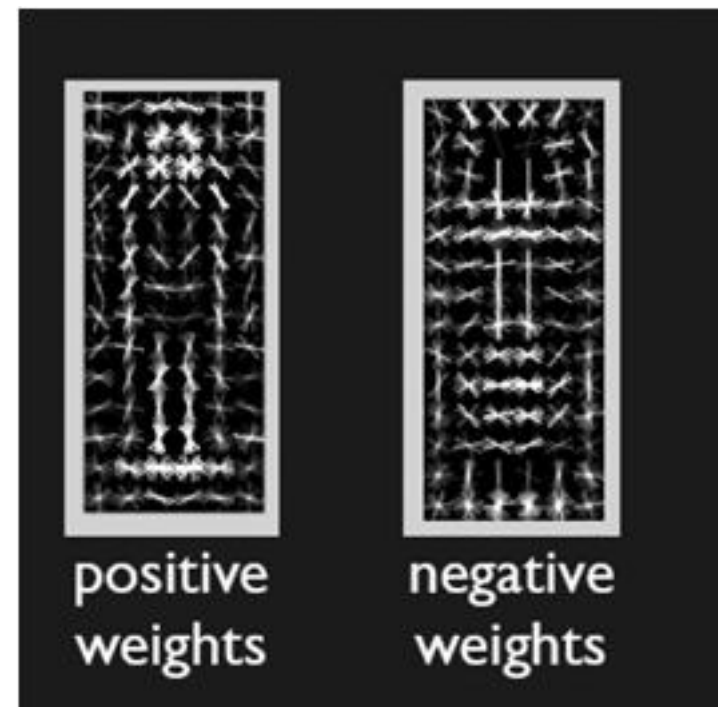
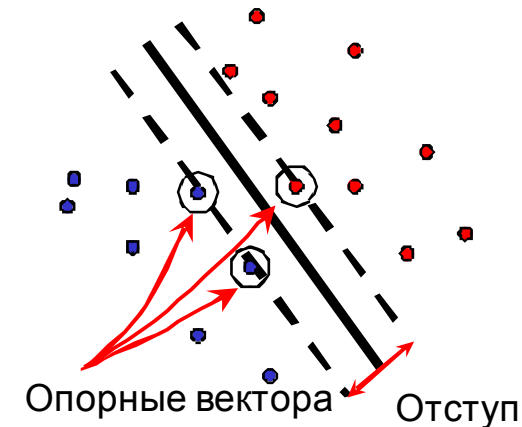
- Опорные вектора с положительными и отрицательными весами
- Чем фактически в нашем случае являются x ?



SVM

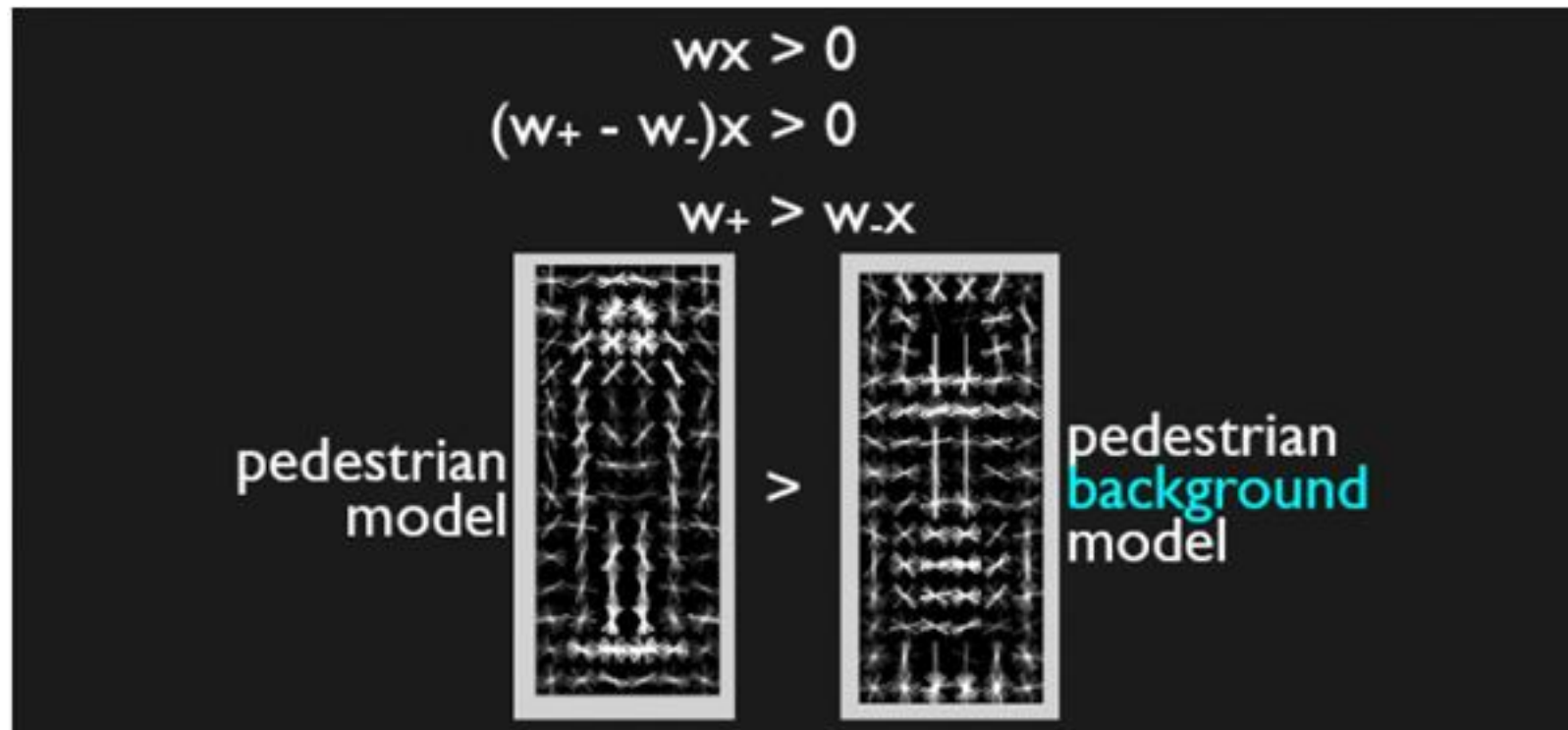


- Каждый опорный вектор – это вектор-признак одного из примеров
- Опорные вектора с положительными и отрицательными весами
- Положительный вес – пример фрагмента, содержащего пешехода
- Отрицательный вес – пример фрагмента, содержащего фон





SVM



- Разделяющая гиперплоскость задается как $w x$ – линейная комбинация положительных и отрицательных примеров
- Каждый признак – «мягкий шаблон» краёв, мы берём не чёткие края, а распределение ориентаций краёв в небольших областях
- Похоже на многократную линейную фильтрацию изображений

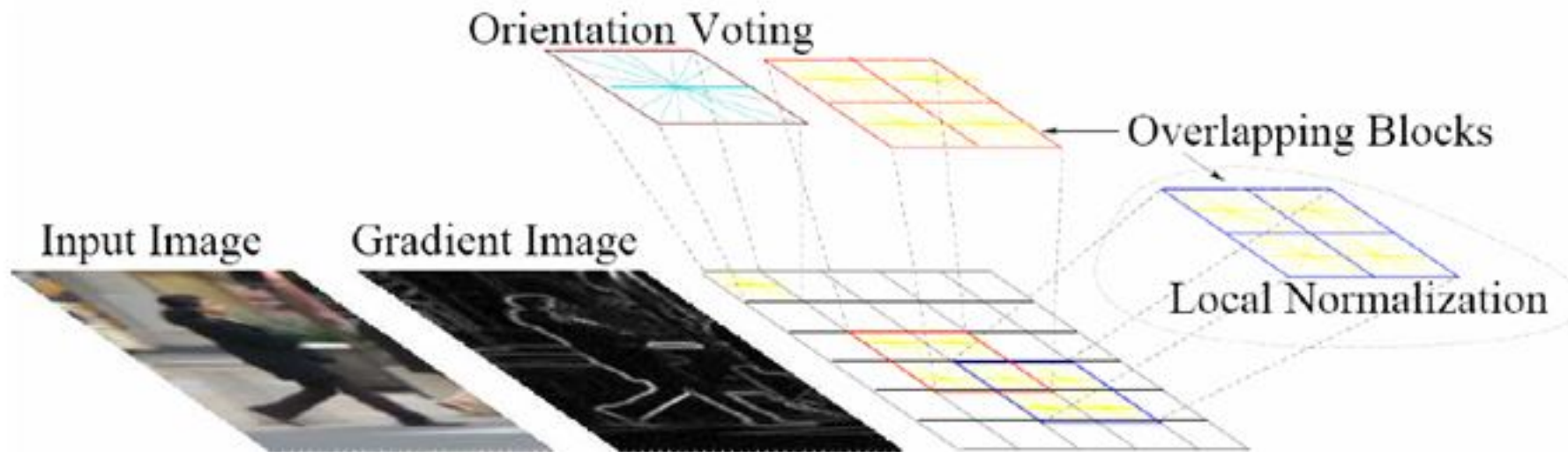


«Детектор пешеходов»

- Получили работающий «детектор пешеходов»
- Алгоритм:
 - Сканирующее окно
 - Вычисление признаков (гистограмм ориентаций градиентов)
 - Классификация с помощью SVM
- Метод работает достаточно неплохо, но неидеально. Надо улучшить.
- В чём проблемы?



Усложнение схемы

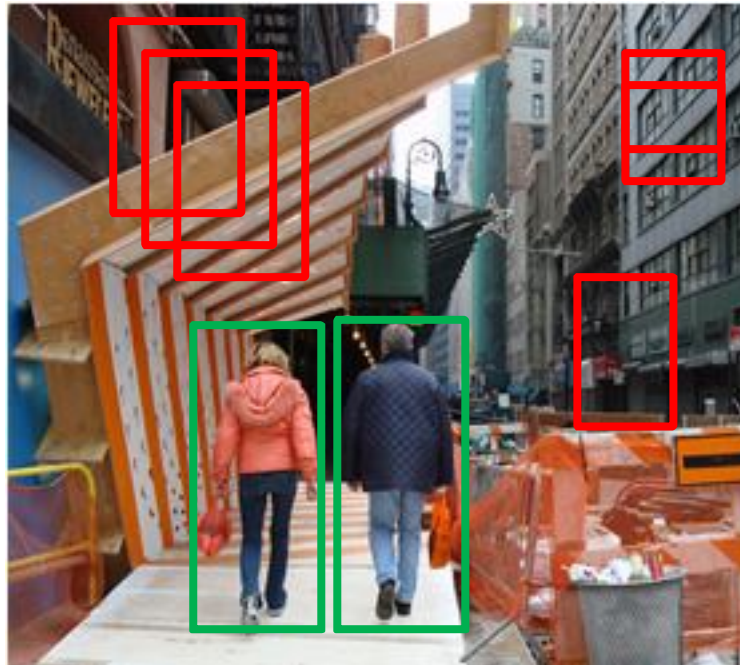


- Добавляем второй уровень:
 - Выделяем блоки 2x2 с перекрытием
 - Нормализуем гистограммы в каждом блоке
 - Это дает большую пространственную поддержку
- Размер вектора = 16 x 8 (сетка) x 8 (ориентаций) x 4 (блоки с перекрытием) = 4096



Обучение детектора

- Сколько на изображении объектов «пешеход» и сколько фрагментов фона?



- Выделение объектов ассиметричная задача: объектов гораздо меньше, чем «не-объектов»
- Вдобавок, класс «не объект» очень сложный – нужно много разных данных для обучения
- Для SVM желательно одинаковое количество и фона, и объекта



Пример – поиск «торса»

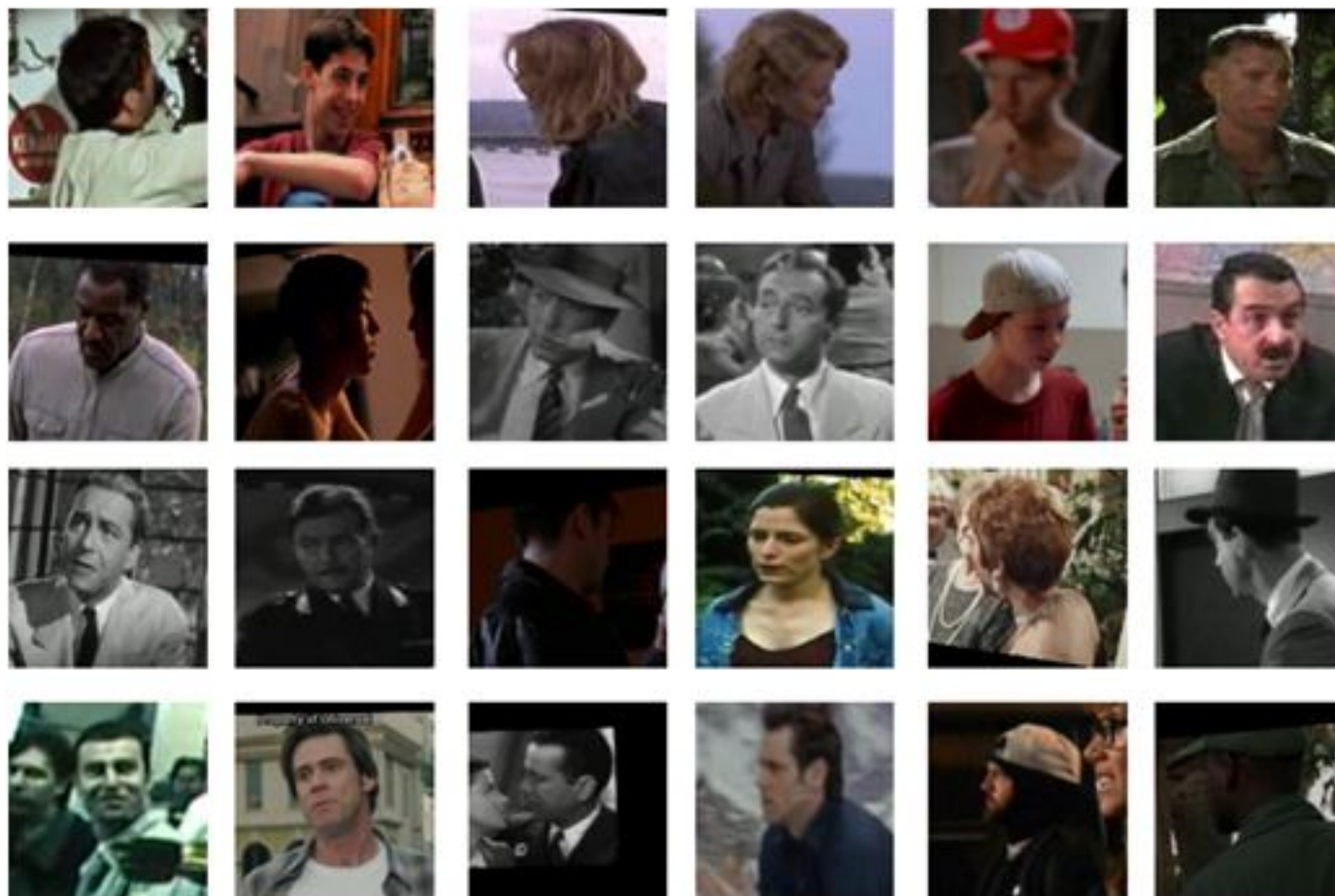
- Хотим построить детектор «верхней части тела и головы»
- Воспользуемся схемой HOG + линейный SVM
- Данные
 - 33 фрагмента фильмов из базы Hollywood2
 - 1122 кадров с размеченными объектами
- На каждом кадре отмечены 1-3 человека, всего 1607 людей, это маловато





Положительные окна

Посмотрим, что отметили люди при разметке:

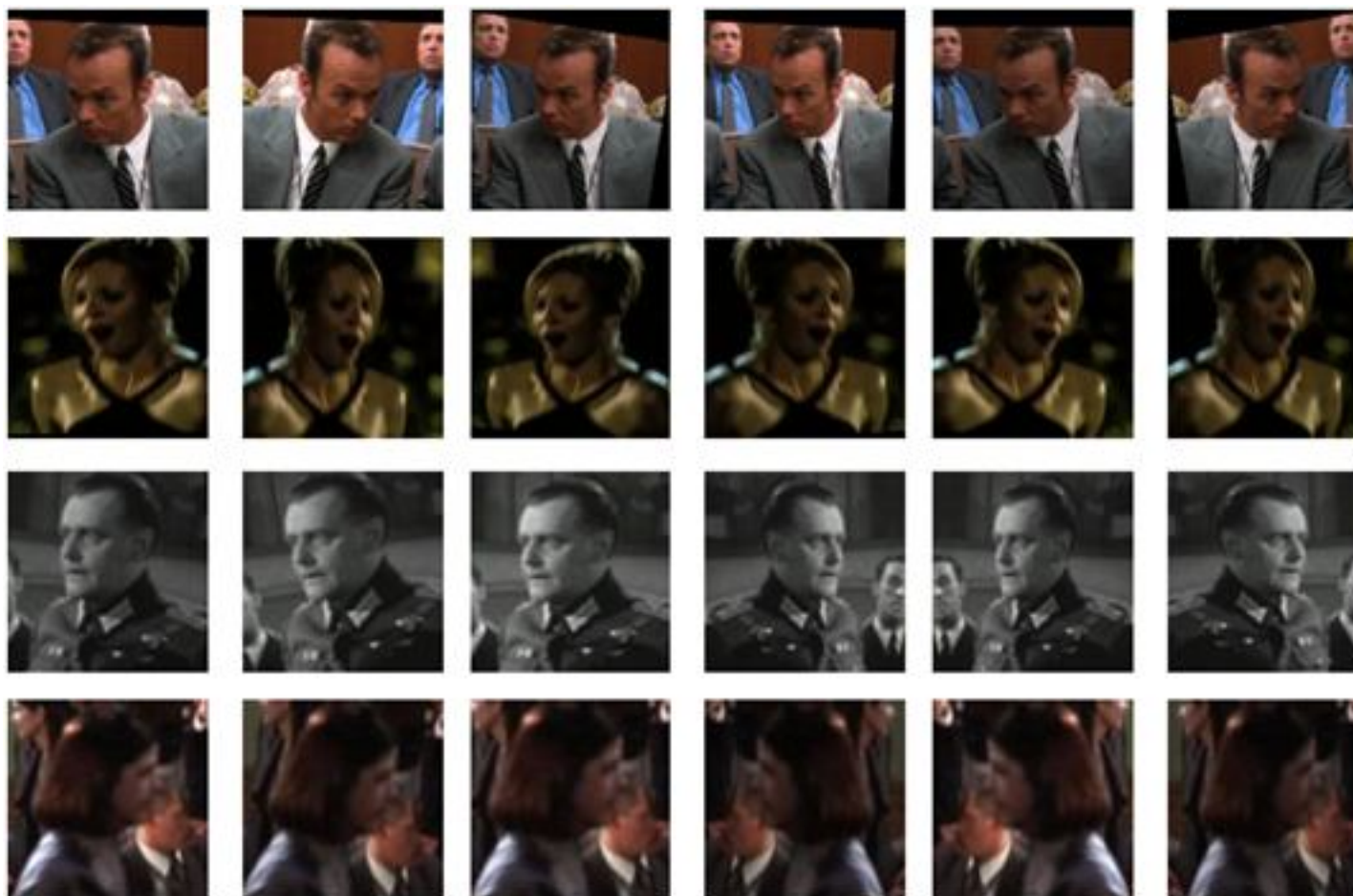


Внимание: похожие положение и ориентация!



Искаженные примеры

Давайте «размножим» данные, «пошевелив» их:



Небольшие сдвиги, отображения, повороты,
изменения масштаба



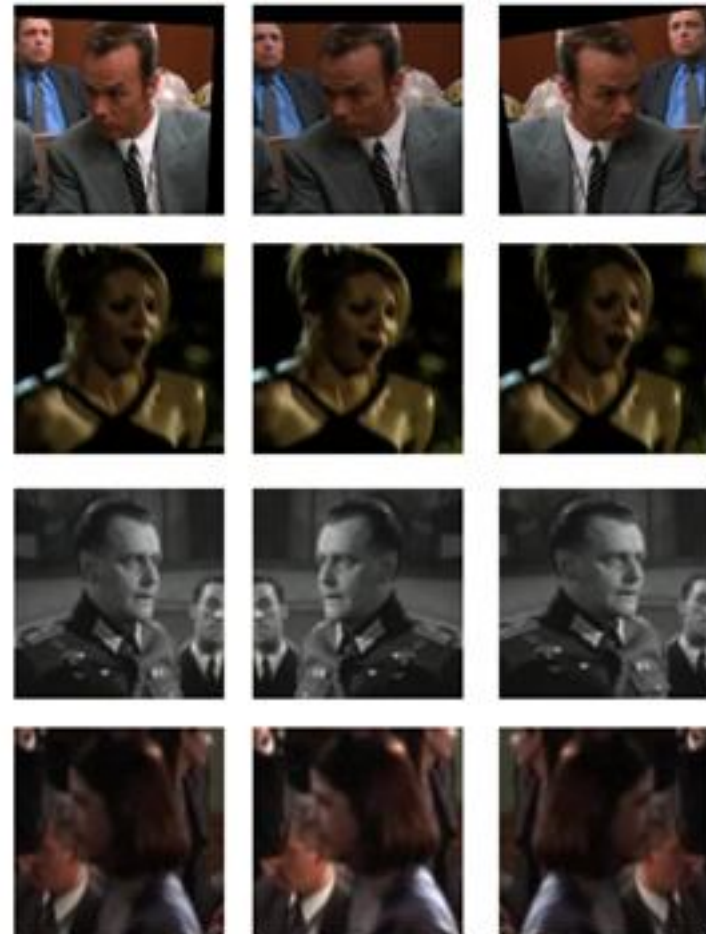
Искаженные примеры

Из 1607 эталонных примеров
получили ~32000 искаженных
(jittered) примеров

Сколько отрицательных примеров
можно набрать из 1100 кадров?

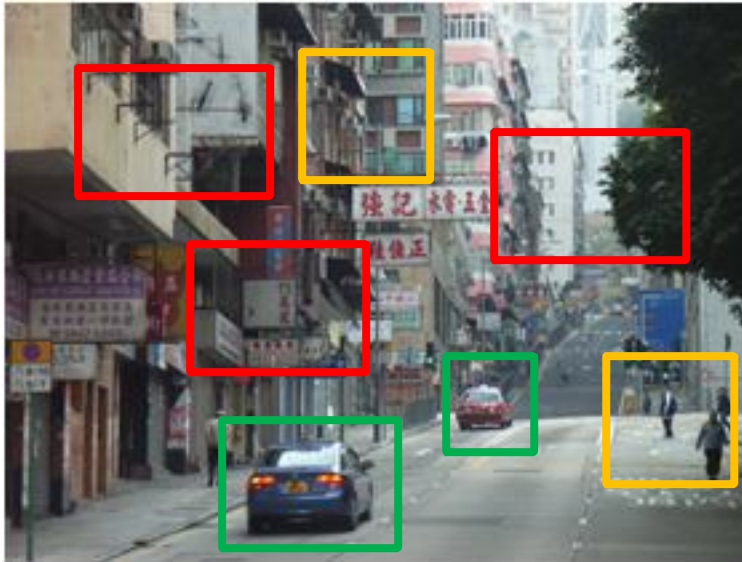
- Гораздо больше 32k.

Вспомним SVM – нам нужны
«трудные примеры» для фона. Как
их найти, если мы всего можем
выбрать ~32k для фона?





Бутстраппинг (Bootstrapping)



- Выбираем отрицательные примеры случайным образом
 - Обучаем классификатор
 - Применяем к данным
 - Добавляем ложные обнаружение к выборке
 - Повторяем
- Смысл:
 - Ложные обнаружения для первого детектора – сложные (**hard negative**)
 - Пусть наша выборка фона будет маленькой, но сложной и представительной



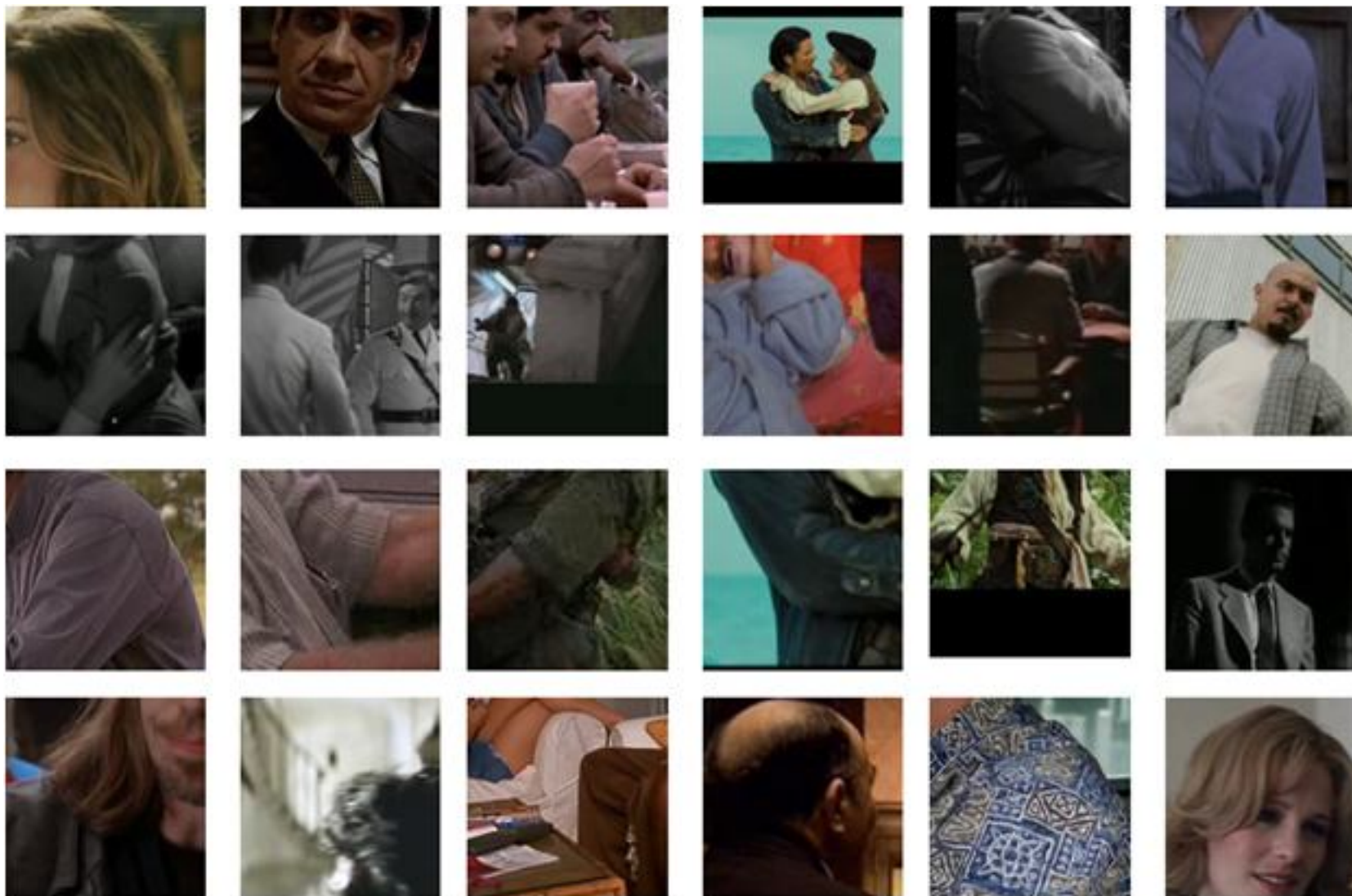
Первая стадия

Трудный
отрицательный
пример

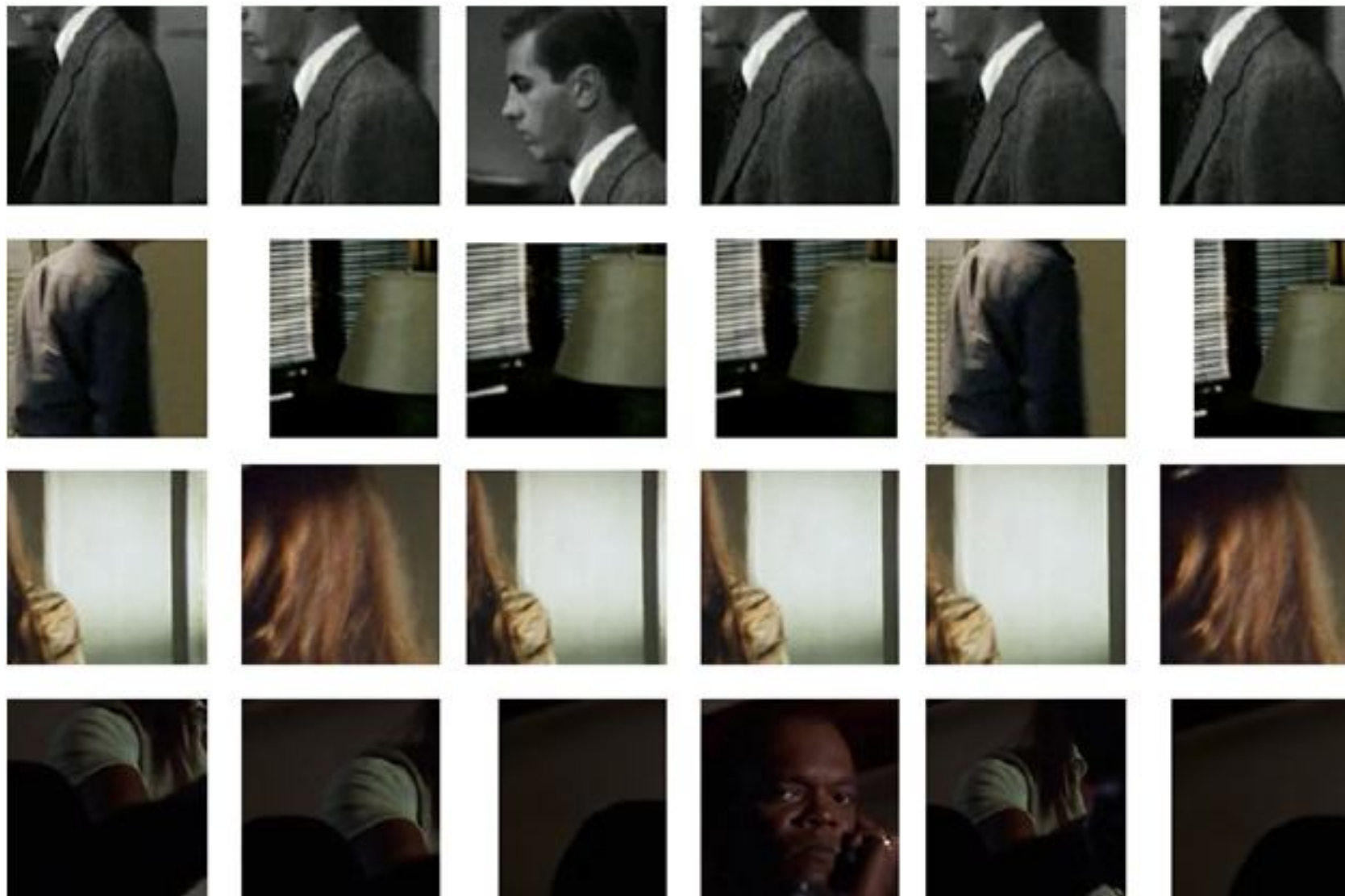


- Ищем ложные обнаружения с высоким рейтингом
- Используем их как трудные отрицательные примеры
- Затраты: # количество изображений x поиск в каждом

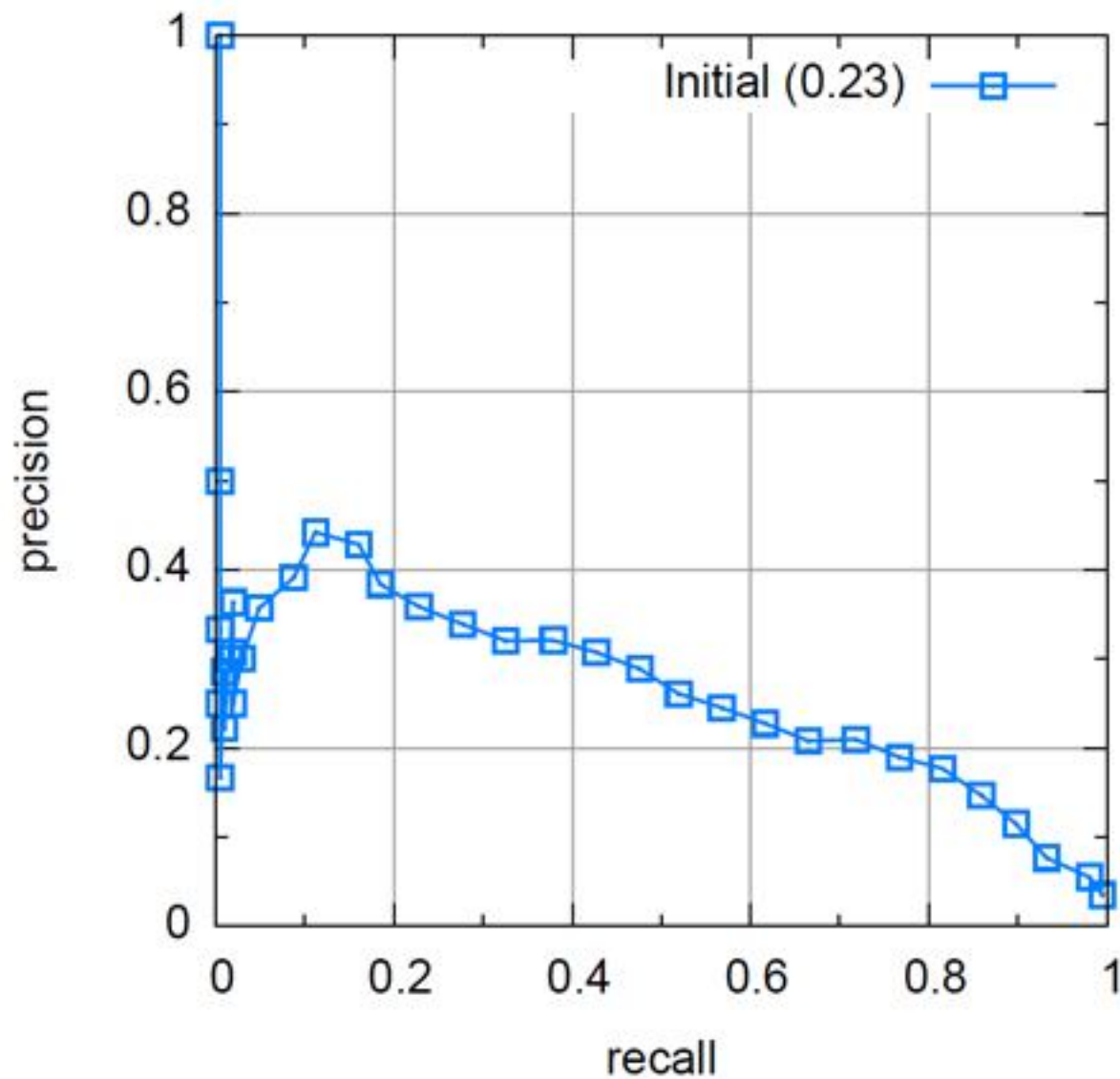
Трудные примеры



Трудные примеры

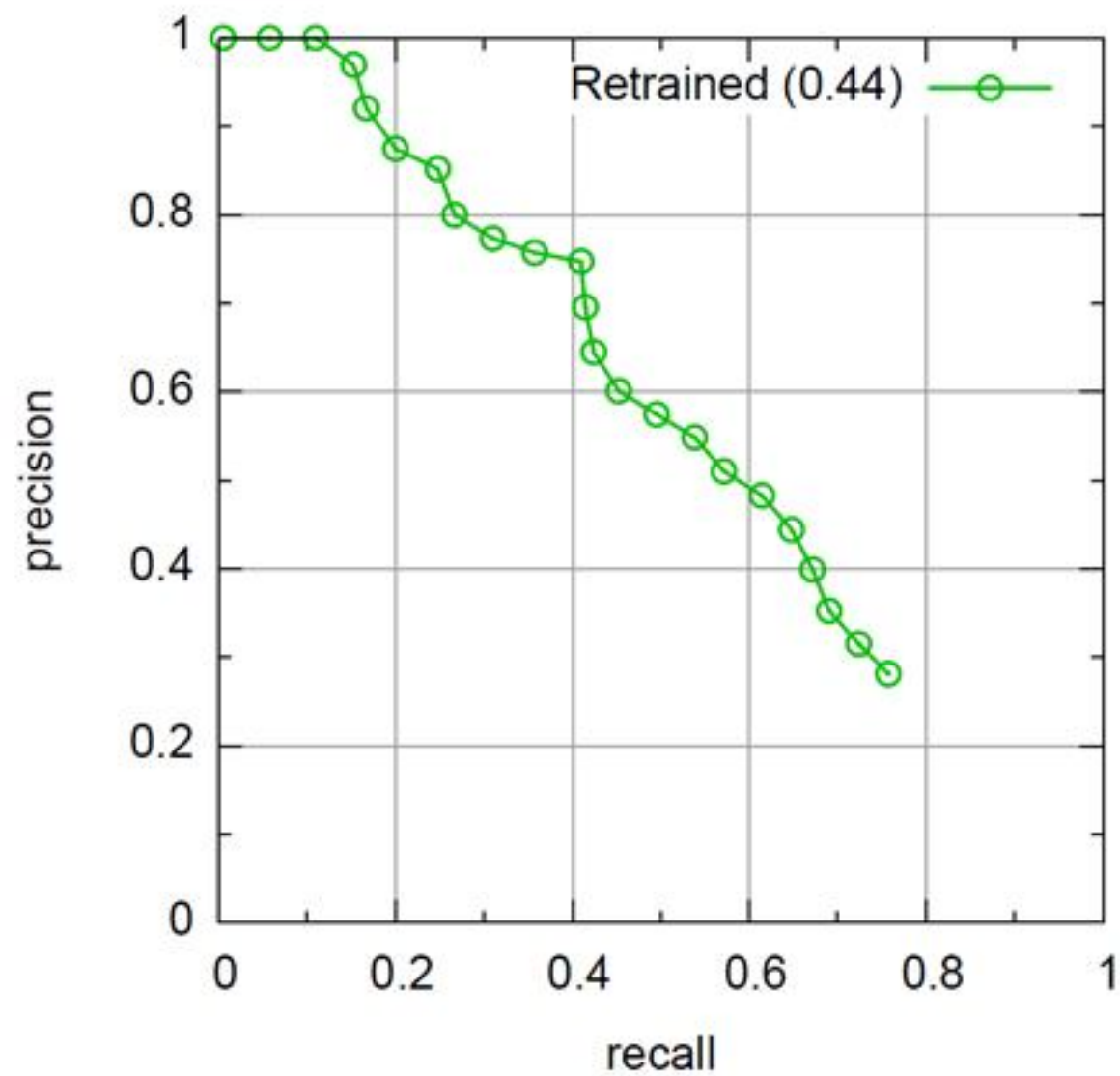


Измерение качества

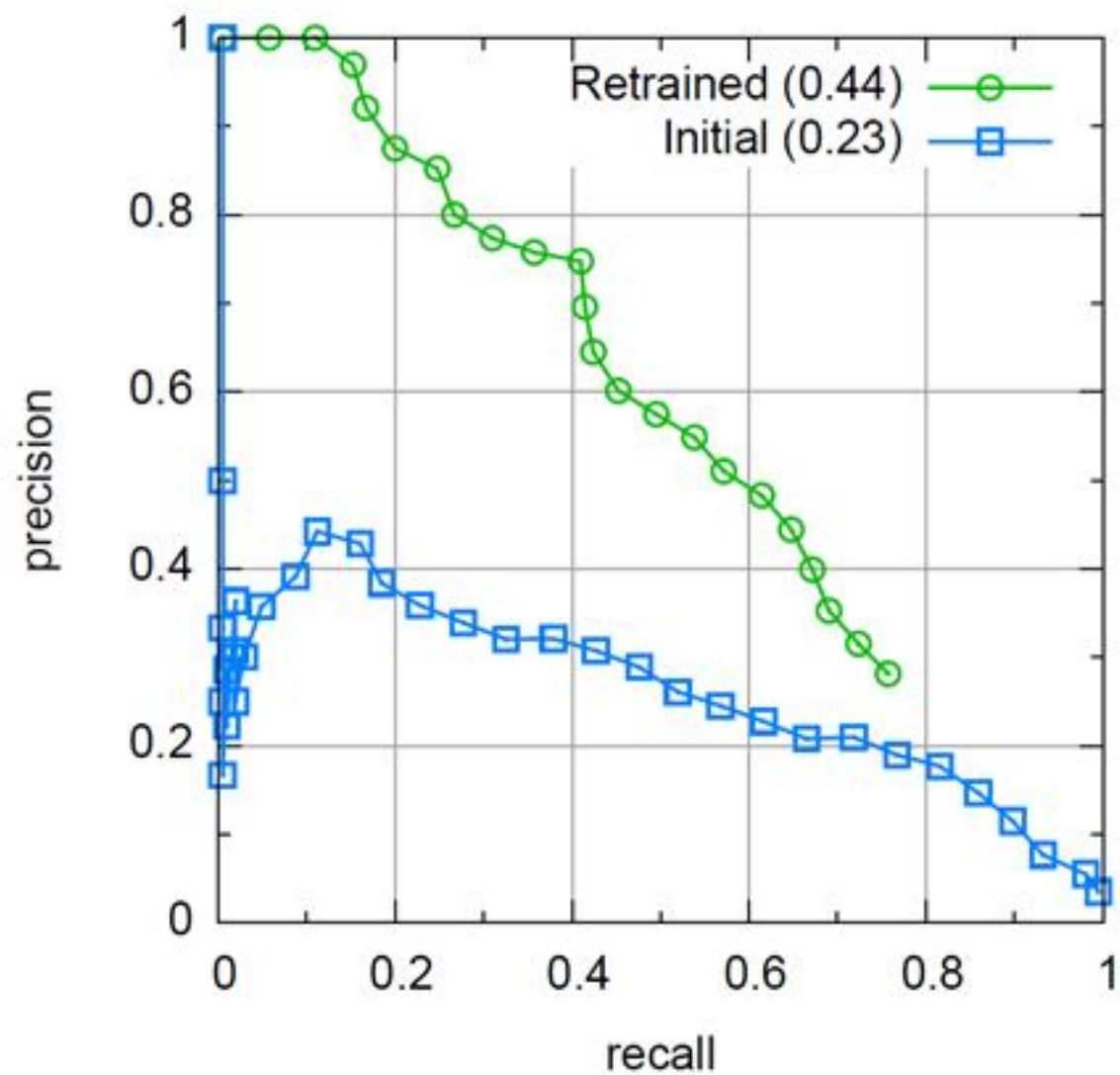




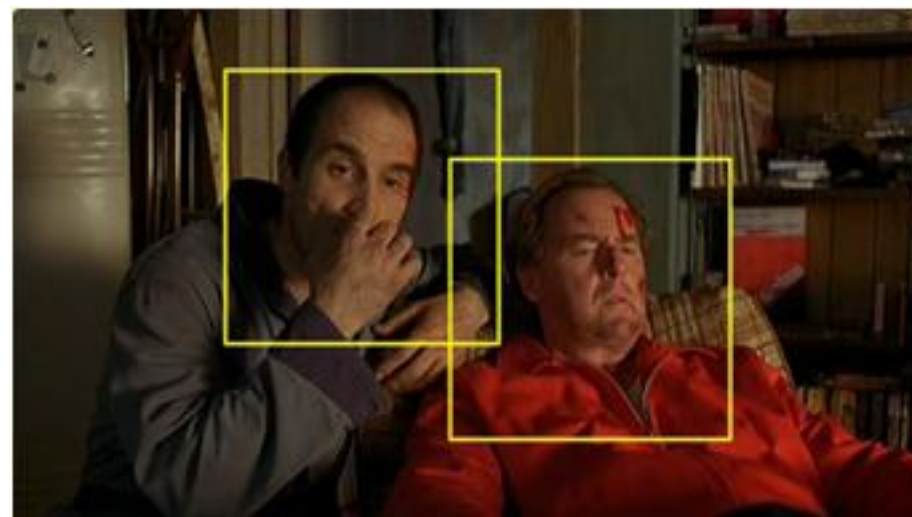
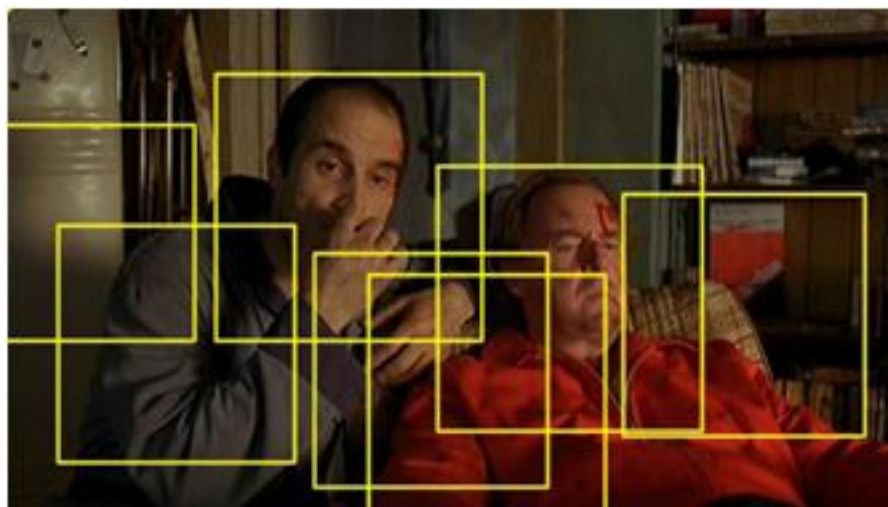
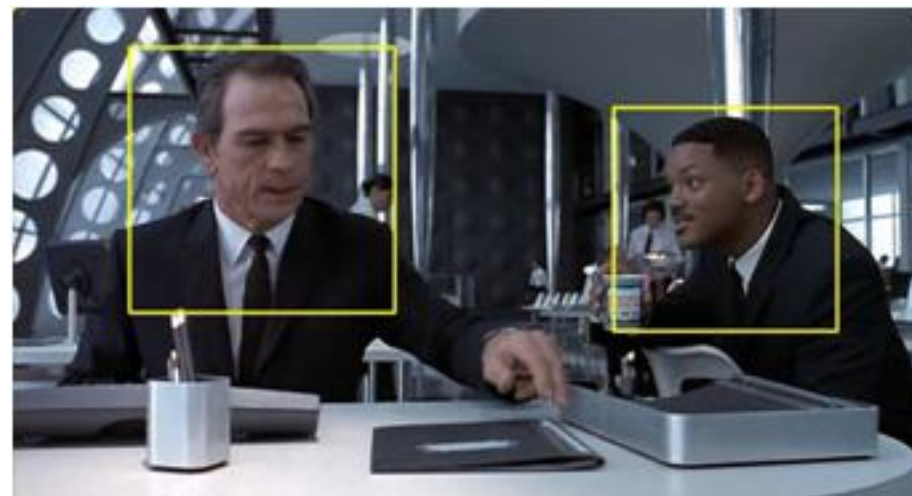
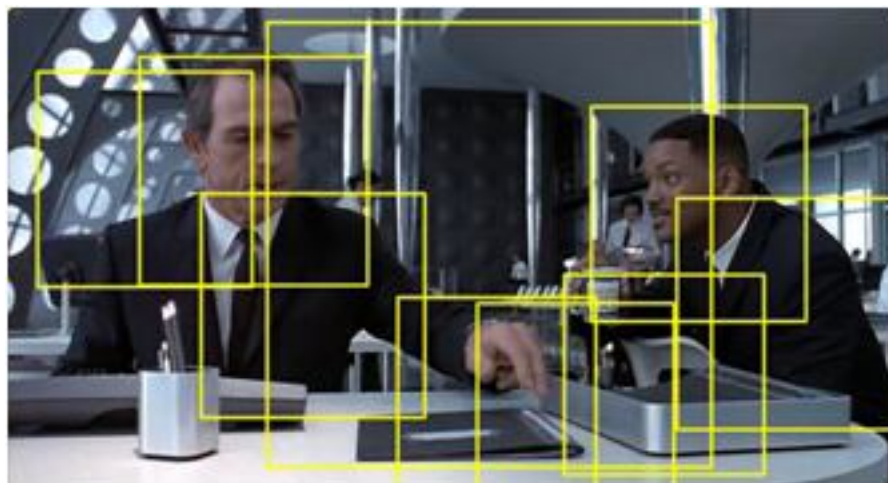
После перетренировки



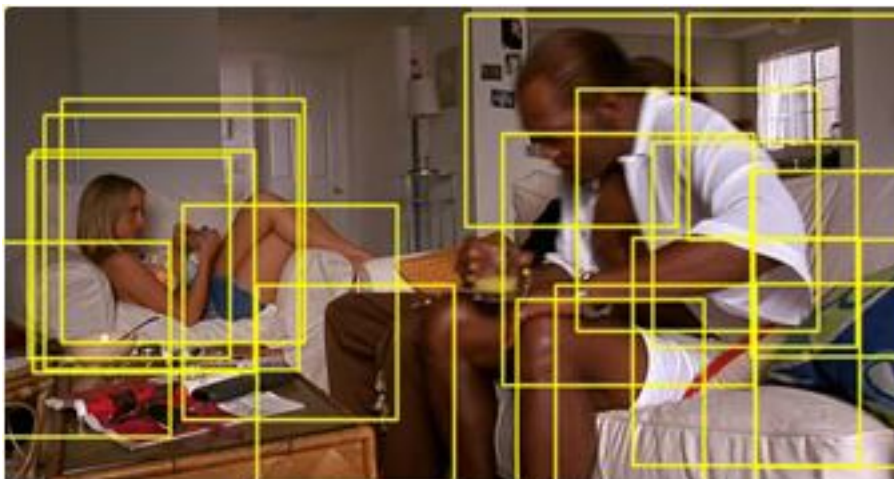
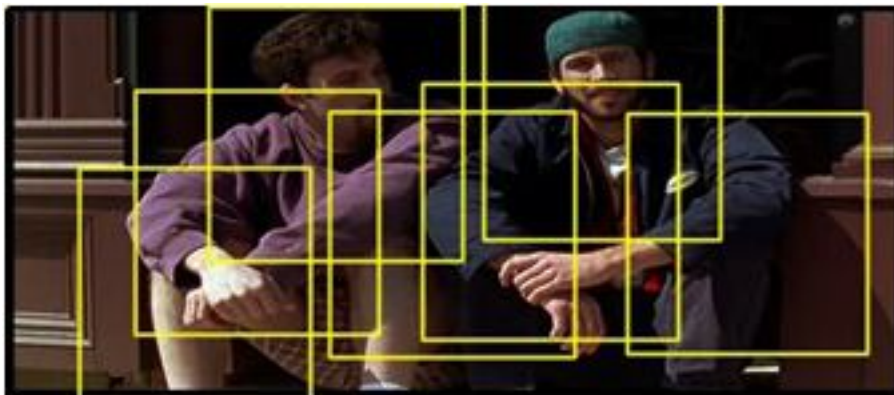
Сравнение



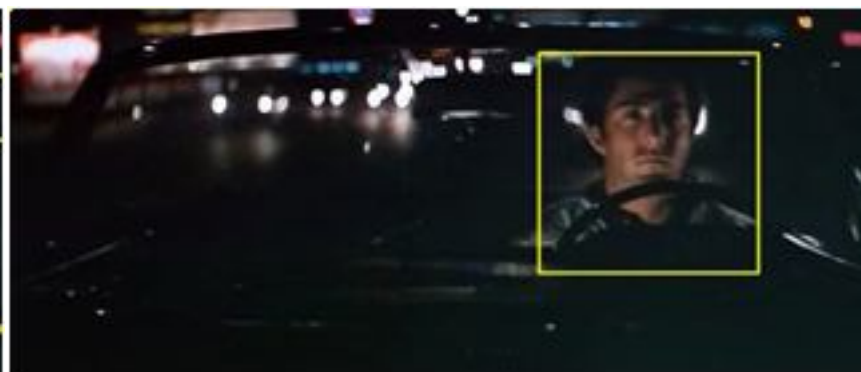
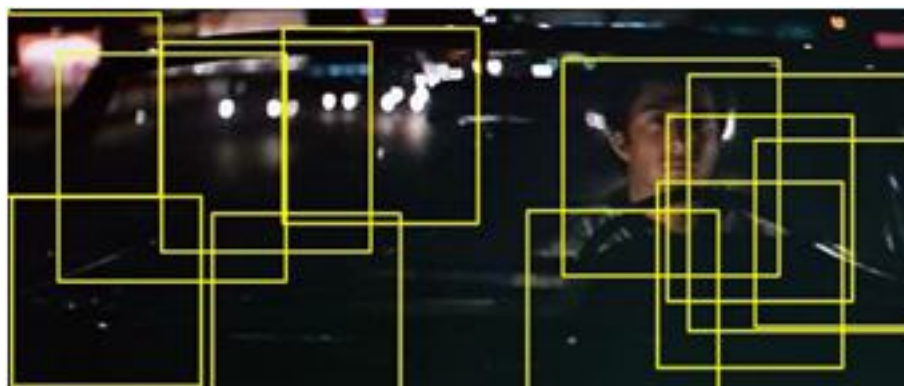
Сравнение



Сравнение



Сравнение





Резюме алгоритма

- Используем скользящее окно
- Вычисляем вектор-признак на основе HOG
 - Разбиваем окно на ячейки
 - В каждой ячейке считаем гистограмму ориентации градиентов
- Обучаемый линейный SVM
- Для обучения:
 - Размножаем (шевелим) эталонные примеры объектов
 - Используем схему bootstrapping для выбора примеров фона
 - На первой стадии берём случайные окна для фона
 - На следующих стадиях выбираем ложные срабатывания детектора как «трудные» примеры



Резюме лекции

- Машинное обучение позволяет использовать большое количество неинтуитивных простых признаков
- Методы и понятия
 - SVM (Метод опорных векторов)
 - HOG (Гистограммы градиентов)
 - Скользящее окно
 - Бутстраппинг и размножение выборки