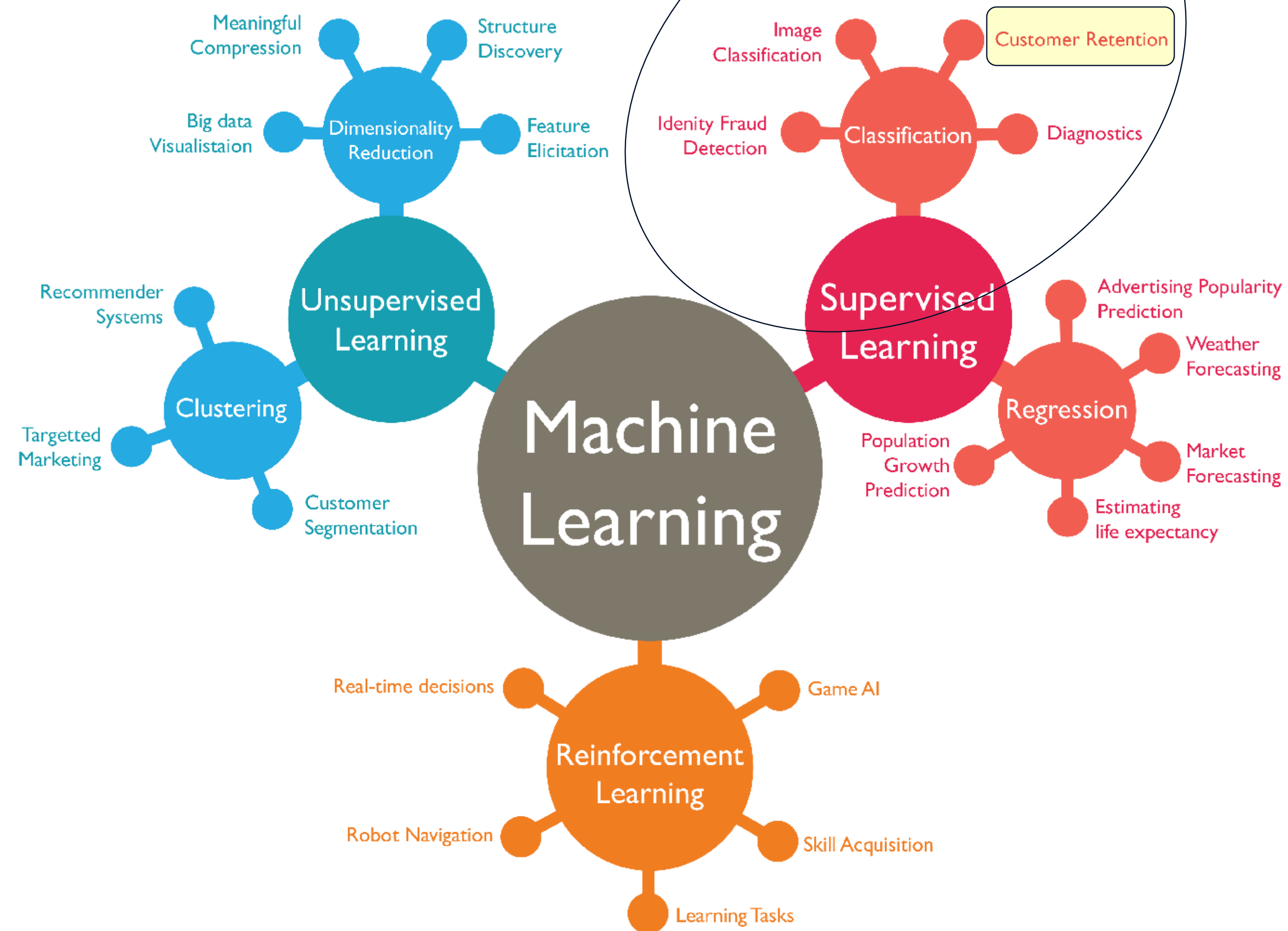


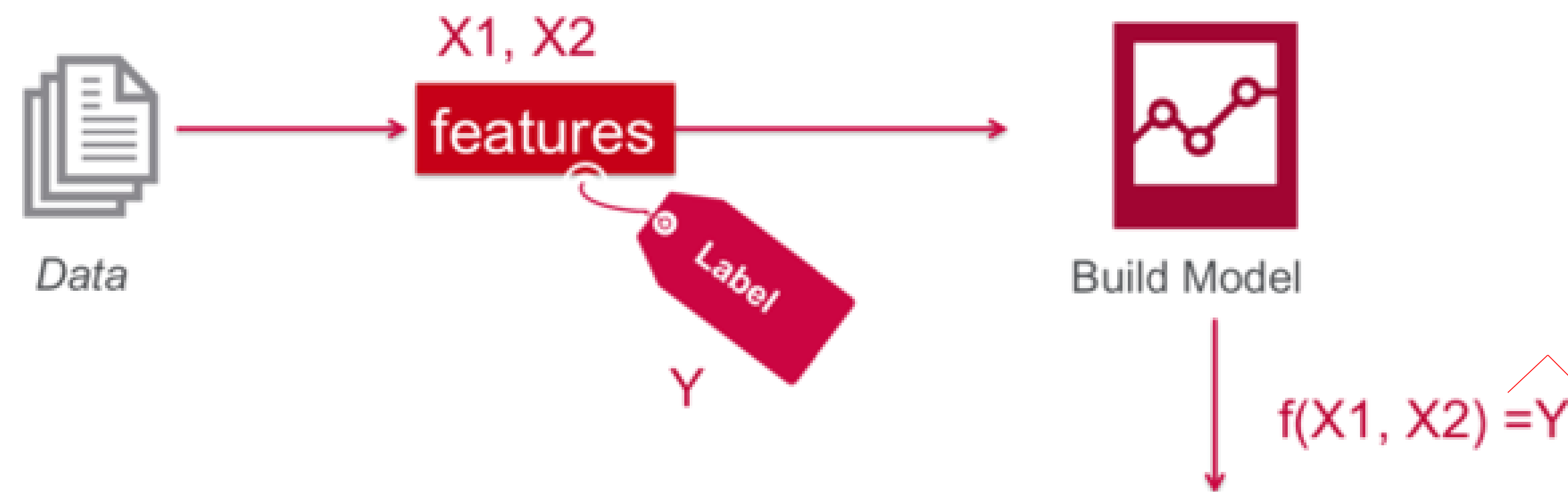


Linear classification

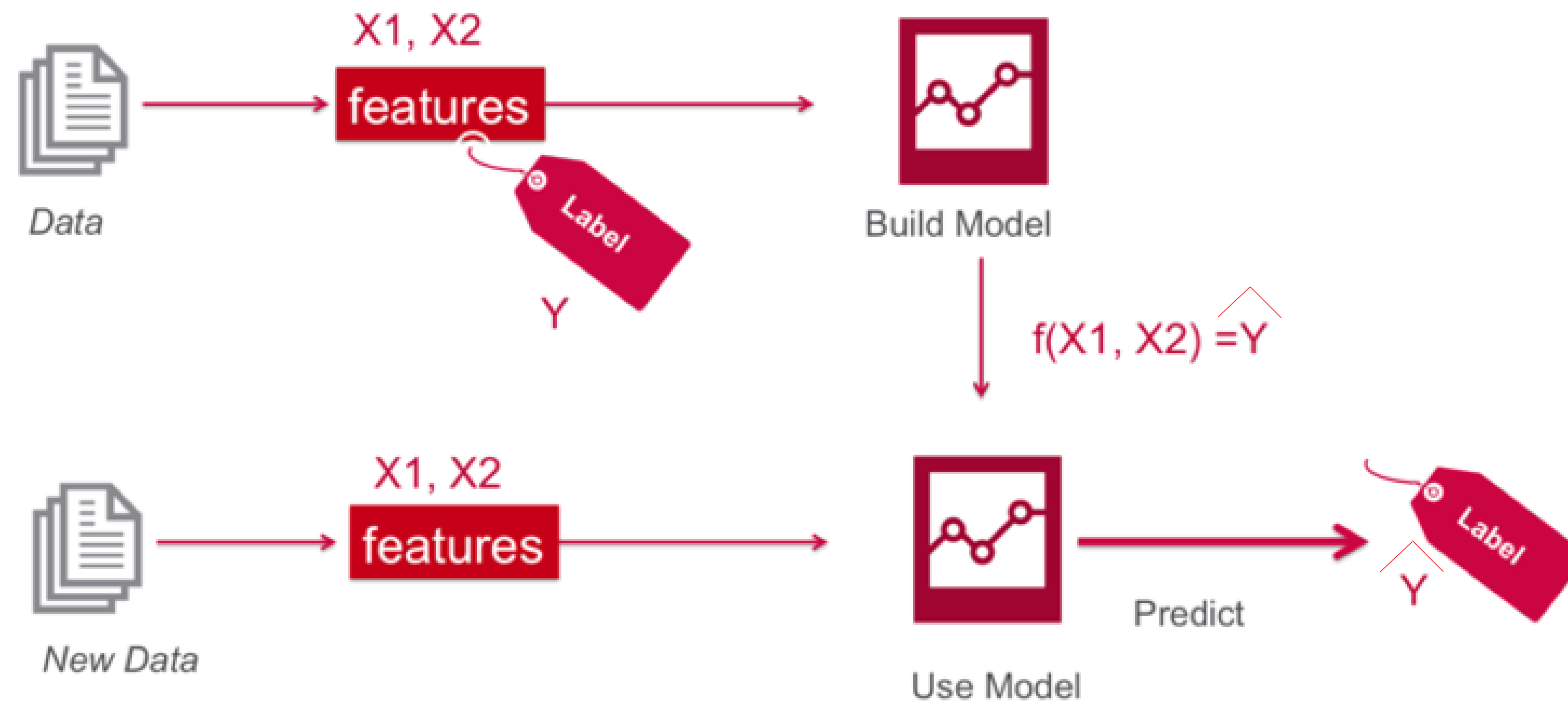
Types of learning



Supervised learning



Supervised learning

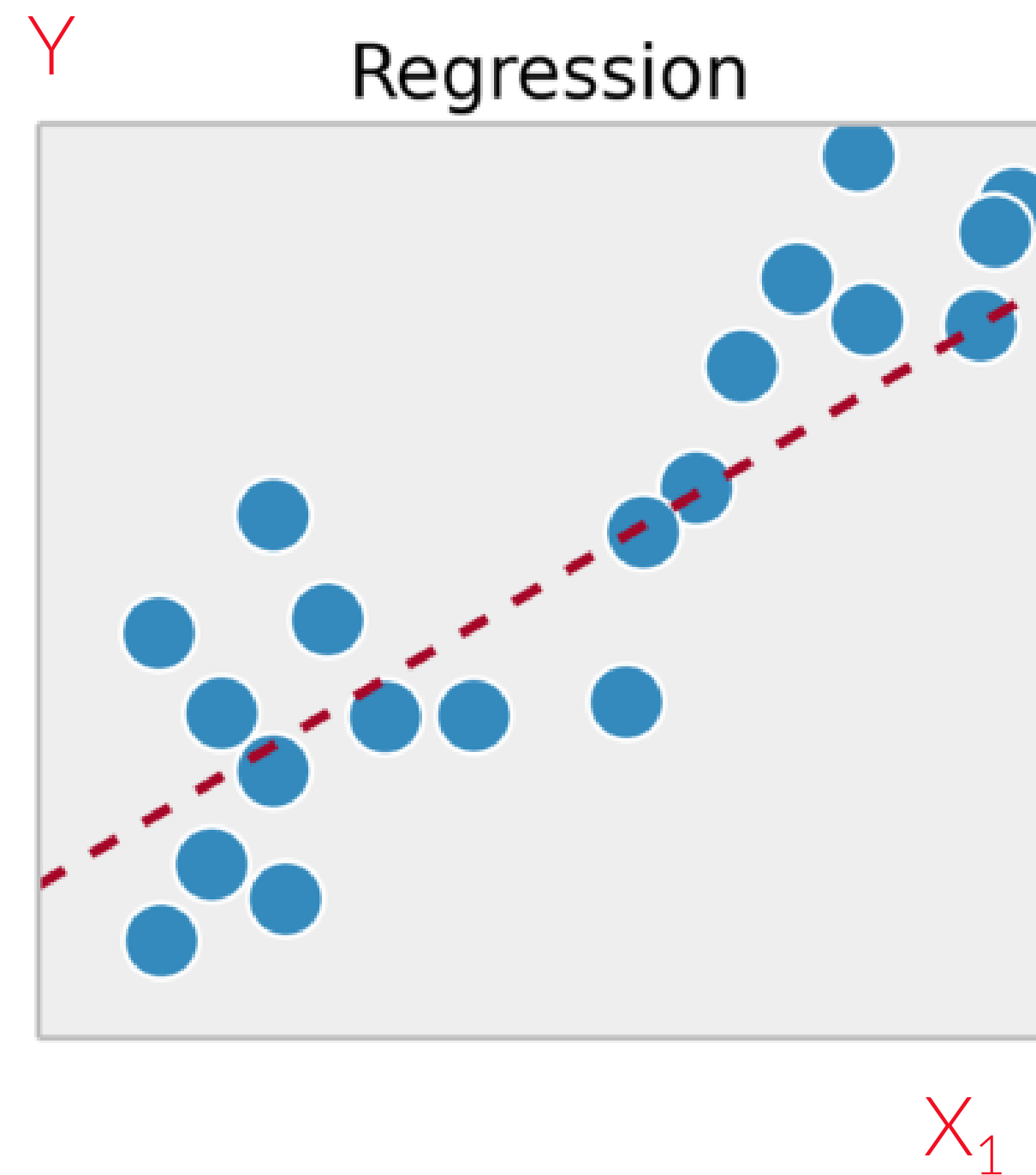


Classification vs Regression

5

Classification: predict a discrete (or categorical) value based on the input variables

Regression: predict a continuous value based on the input variables

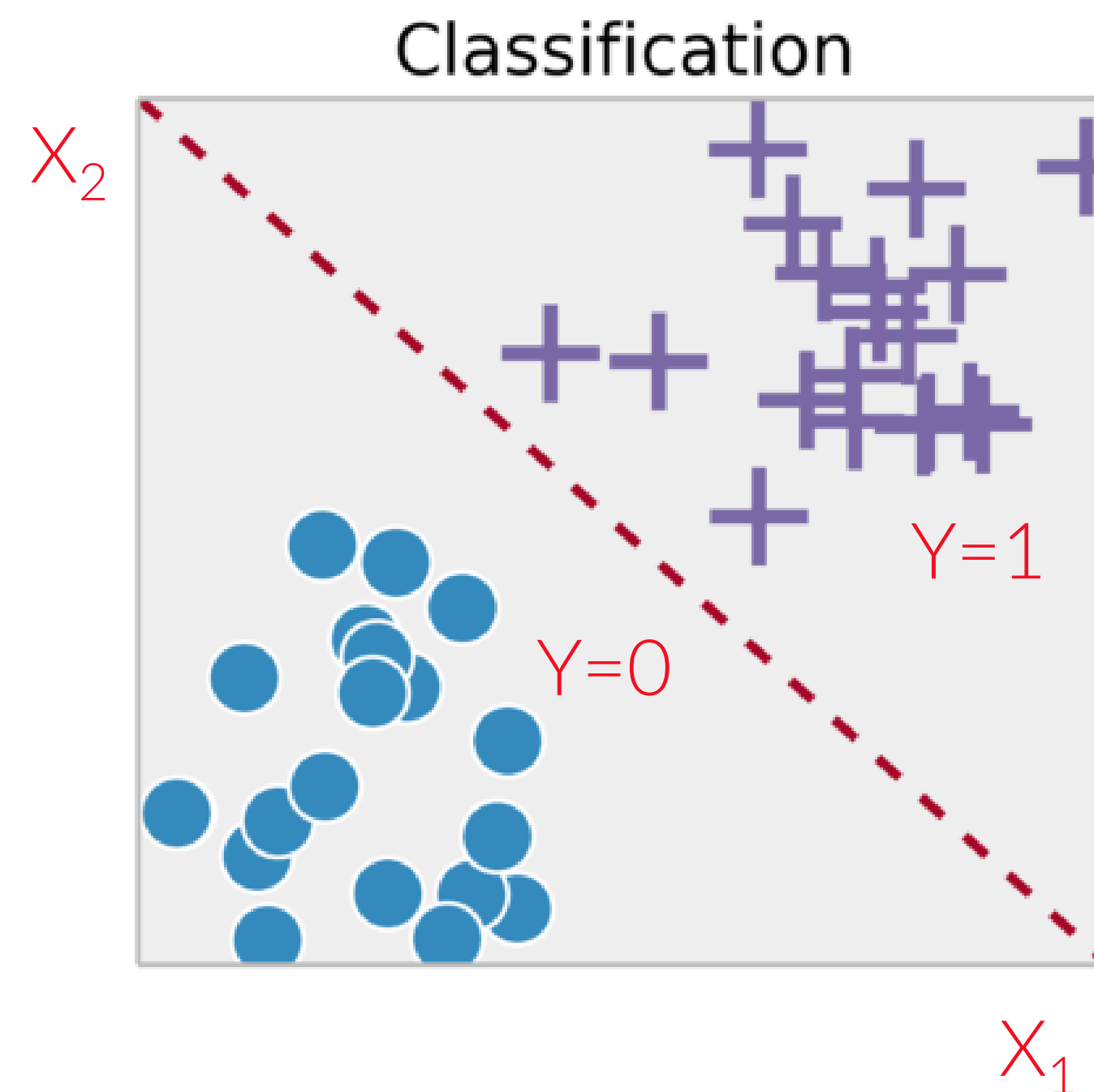


Classification vs Regression

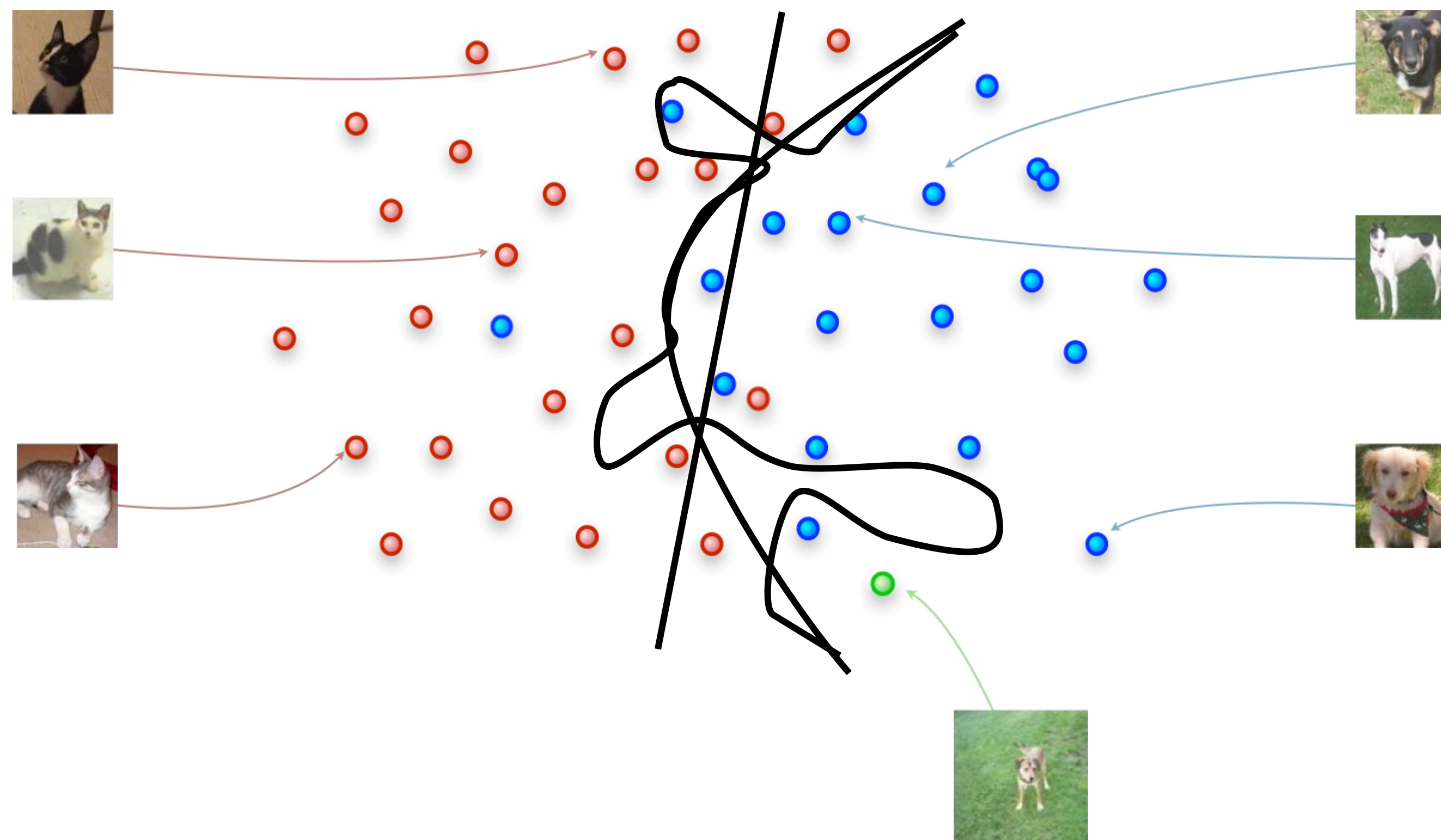
6

Classification: predict a discrete (or categorical) value based on the input variables

Regression: predict a continuous value based on the input variables



Binary classification

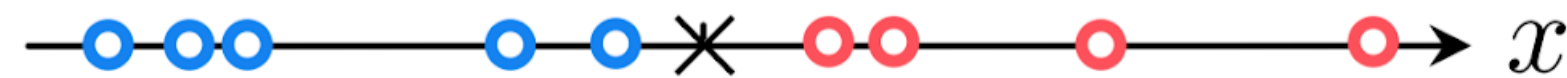


Logistic regression

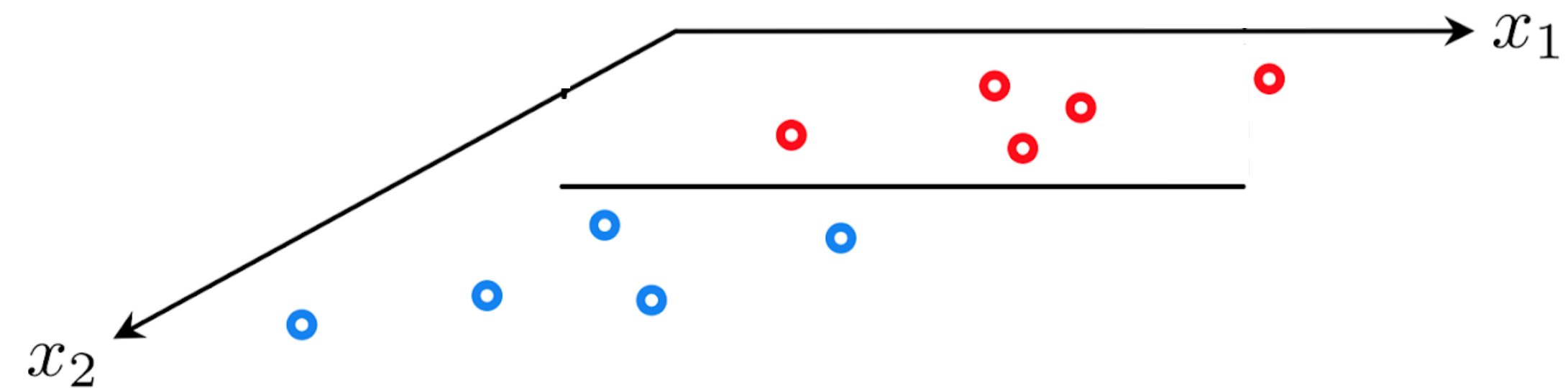
Linear classification

(Focus on binary)

Two perspectives on classification

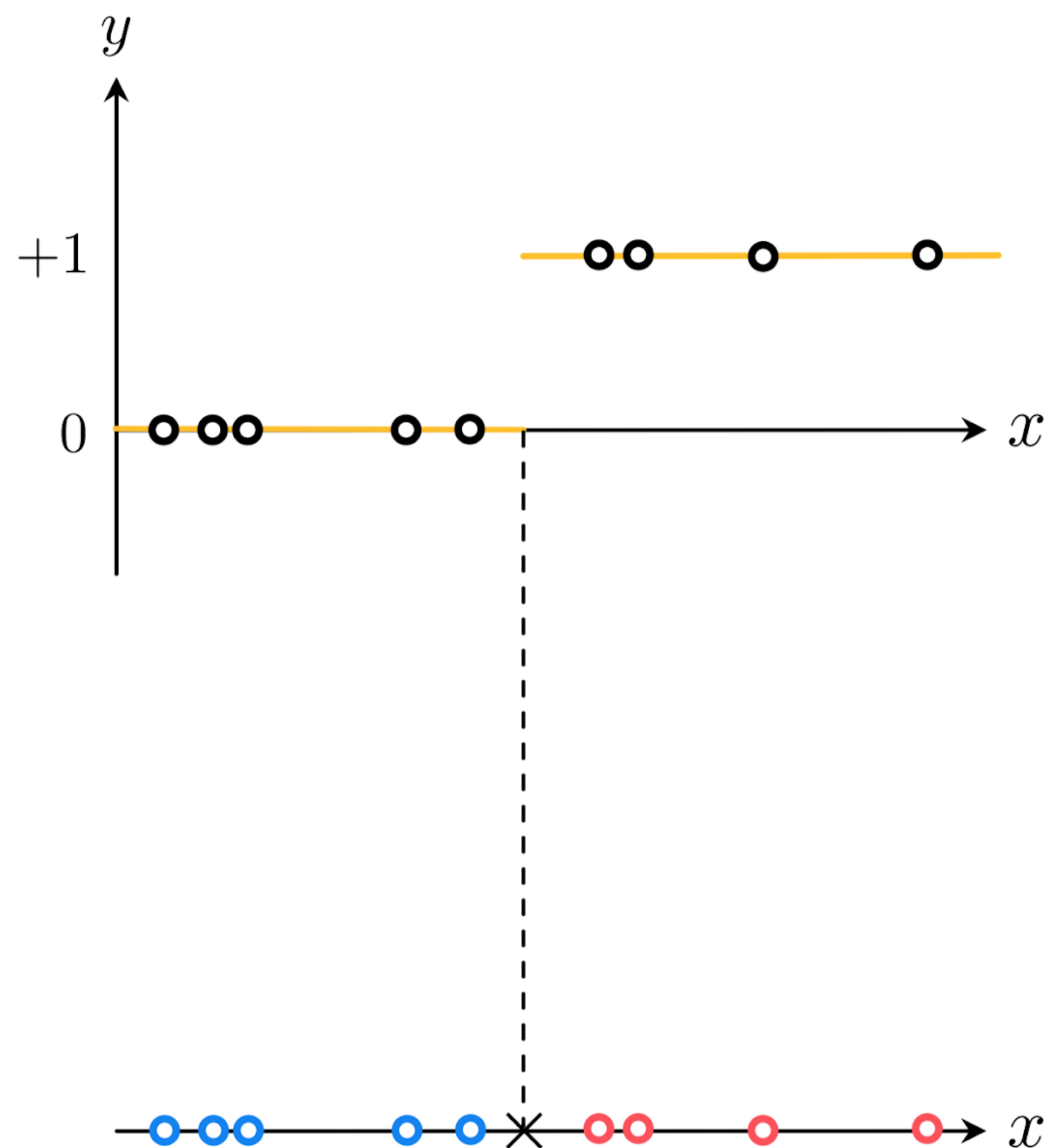


1d input: decision boundary is a point

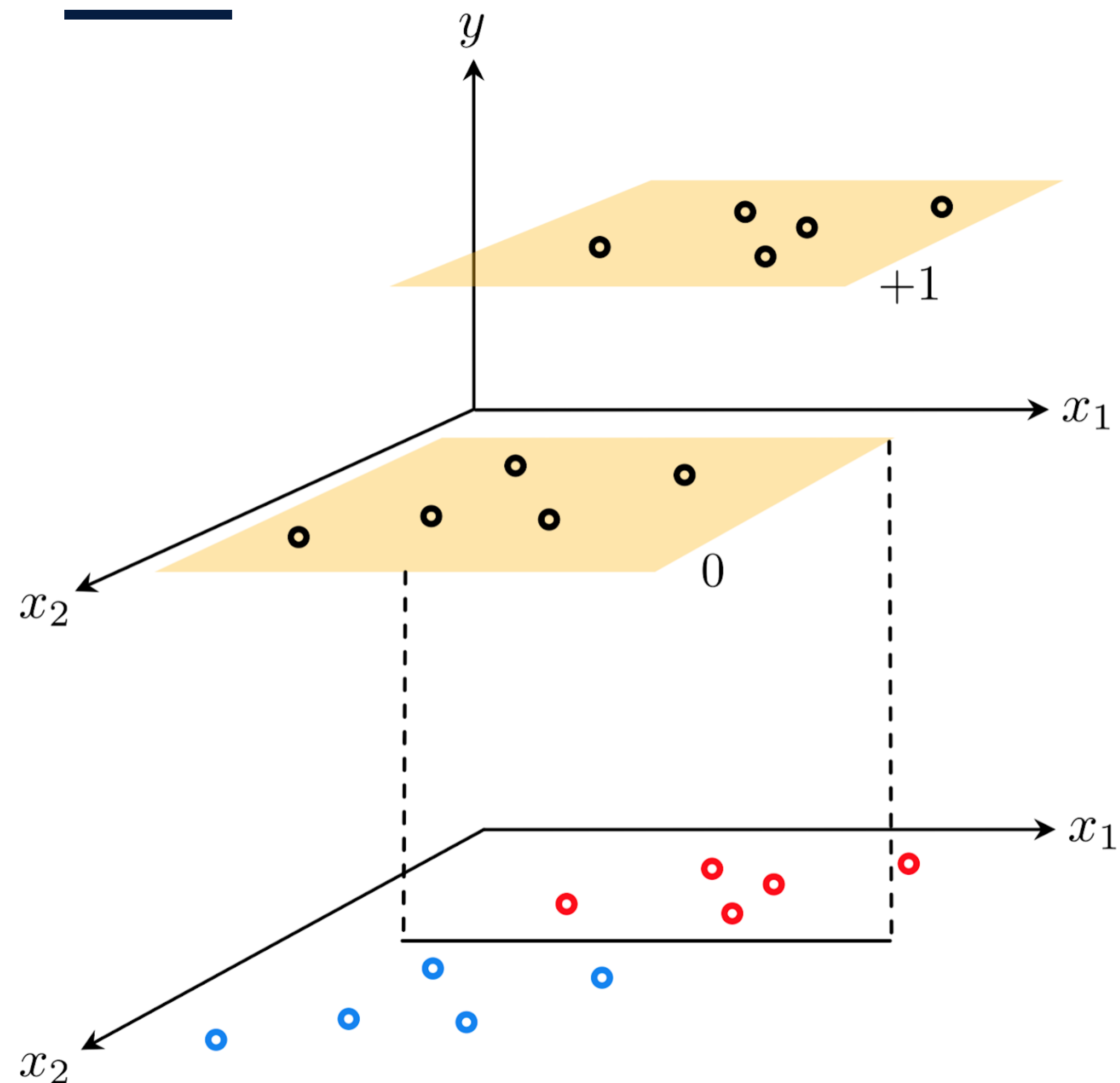


2d input: decision boundary is a line

Two perspectives on classification



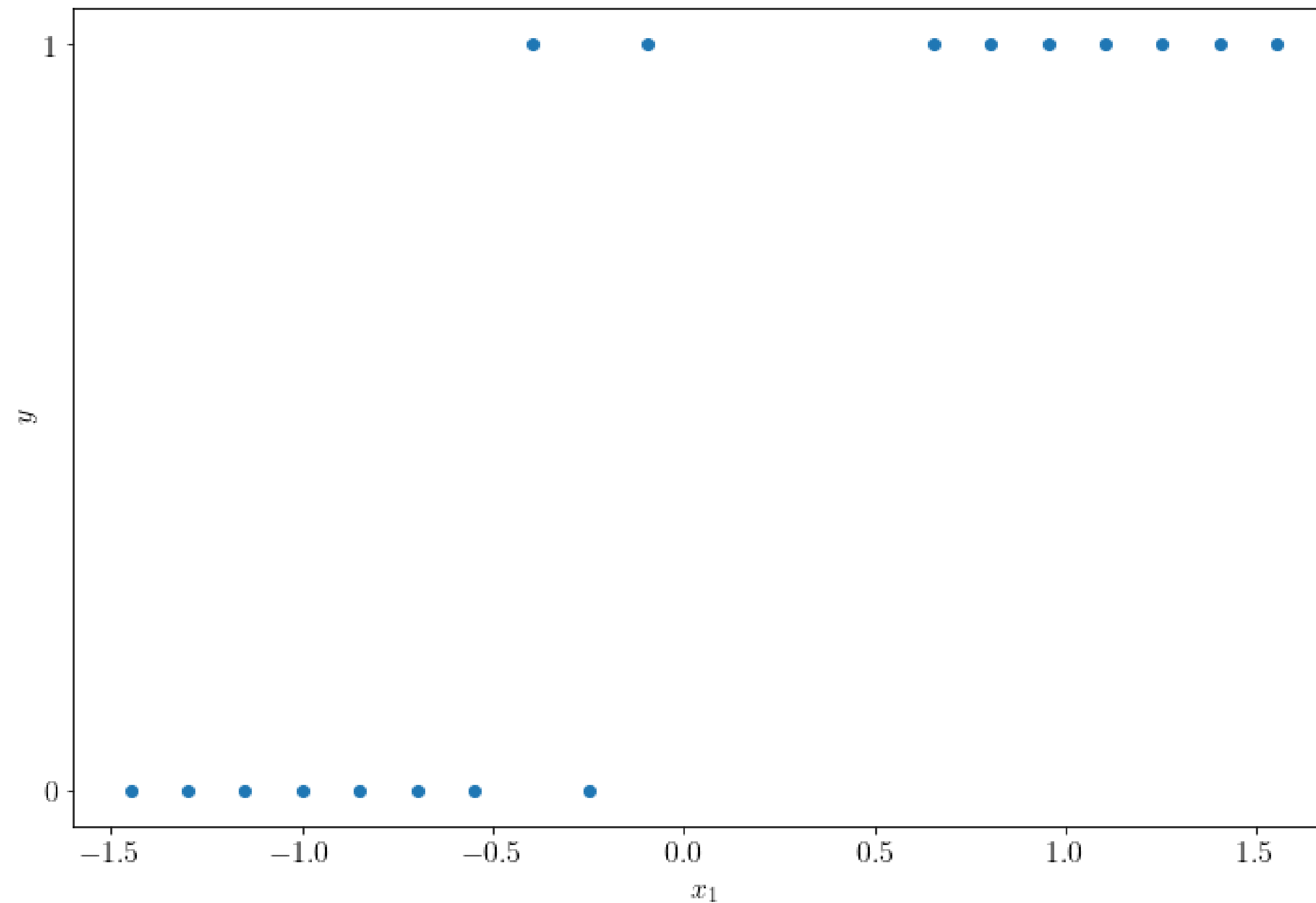
1d input: decision boundary is a point



2d input: decision boundary is a line

Illustrative 1D example

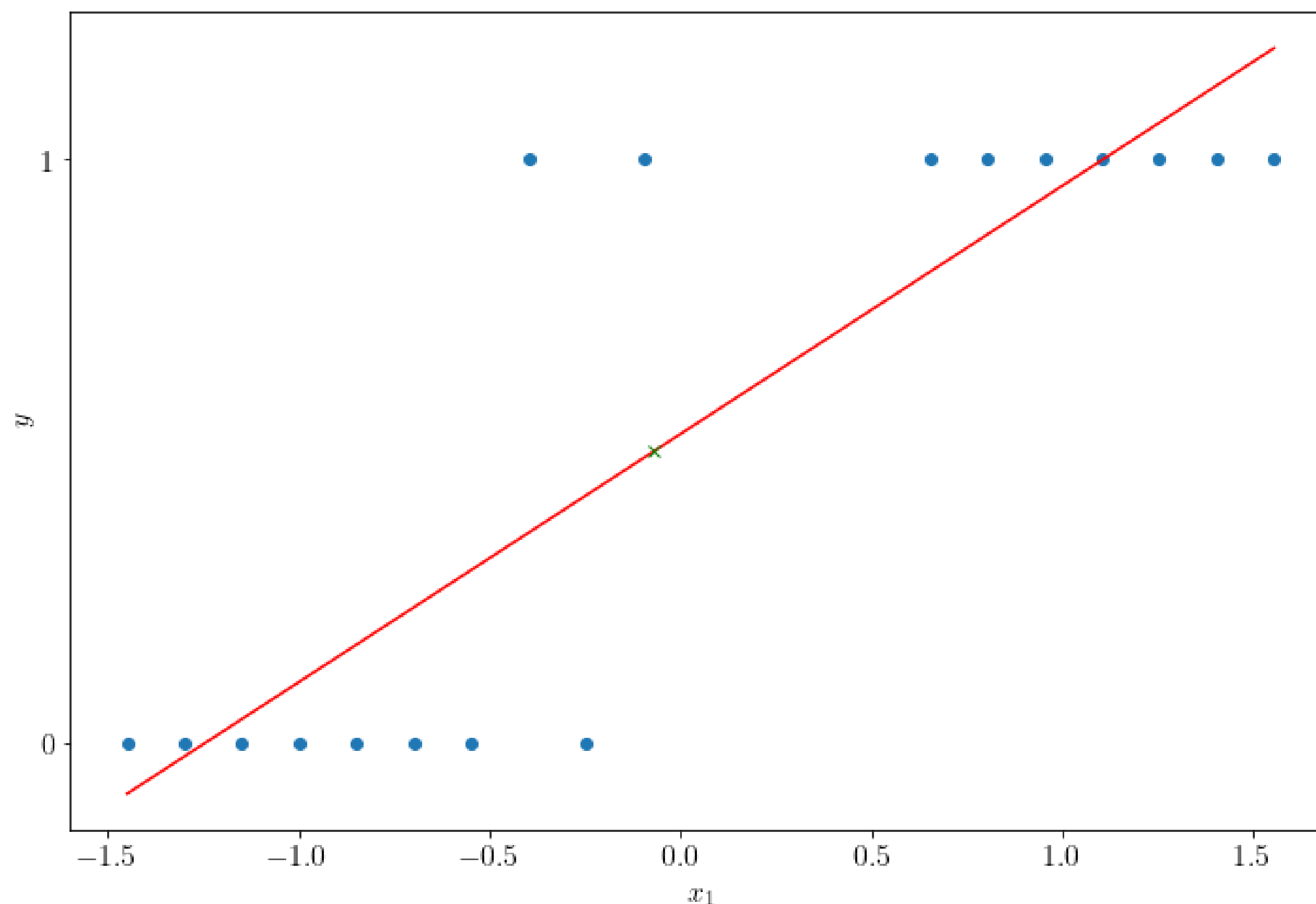
11



Proposition #1

How about applying the linear regression framework seen earlier to fit such data?

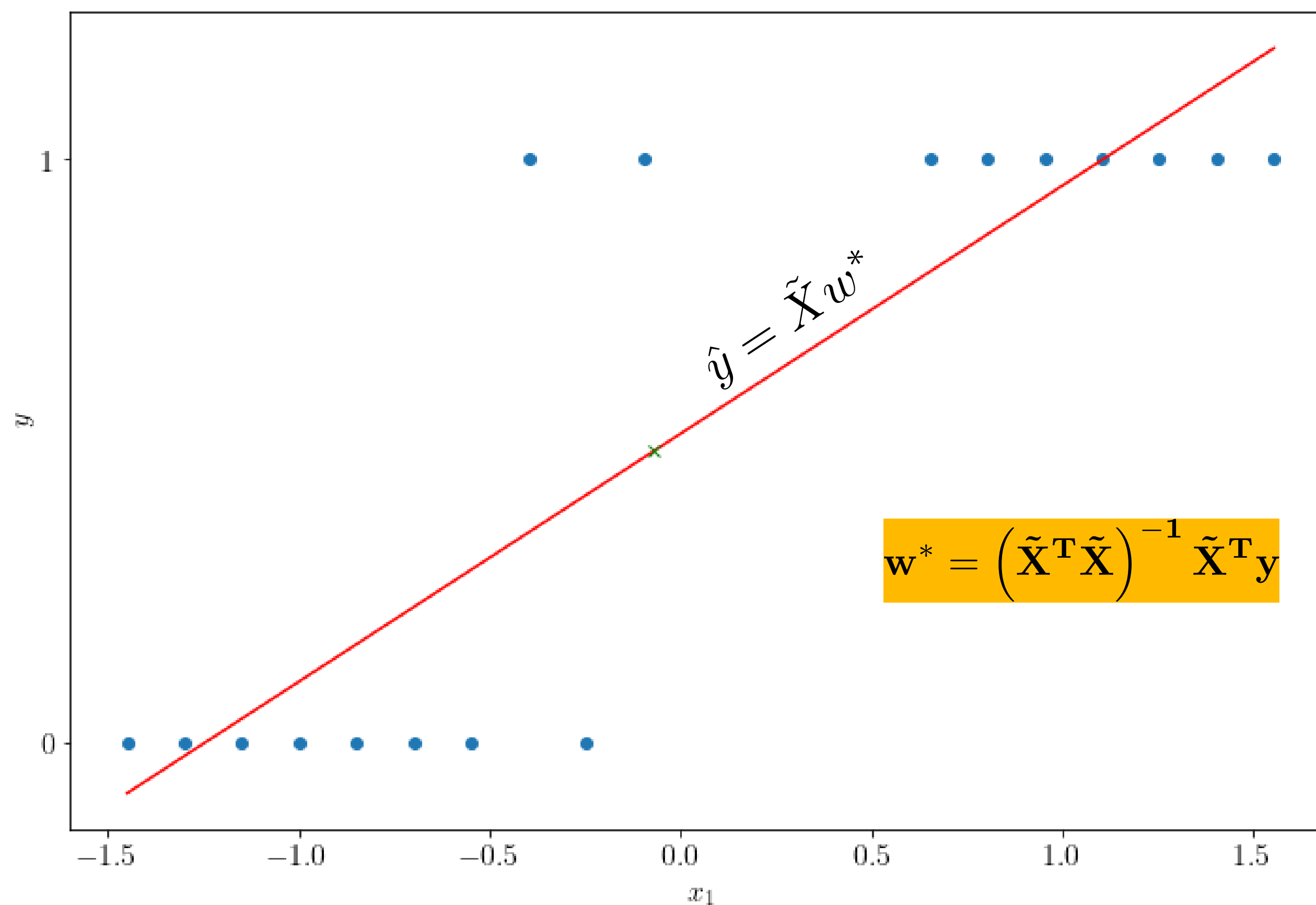
We minimize the following cost function $J(w) = \frac{1}{N} \sum_{i=1}^N \left(\tilde{X}^i w - y^i \right)^2$



Proposition #1

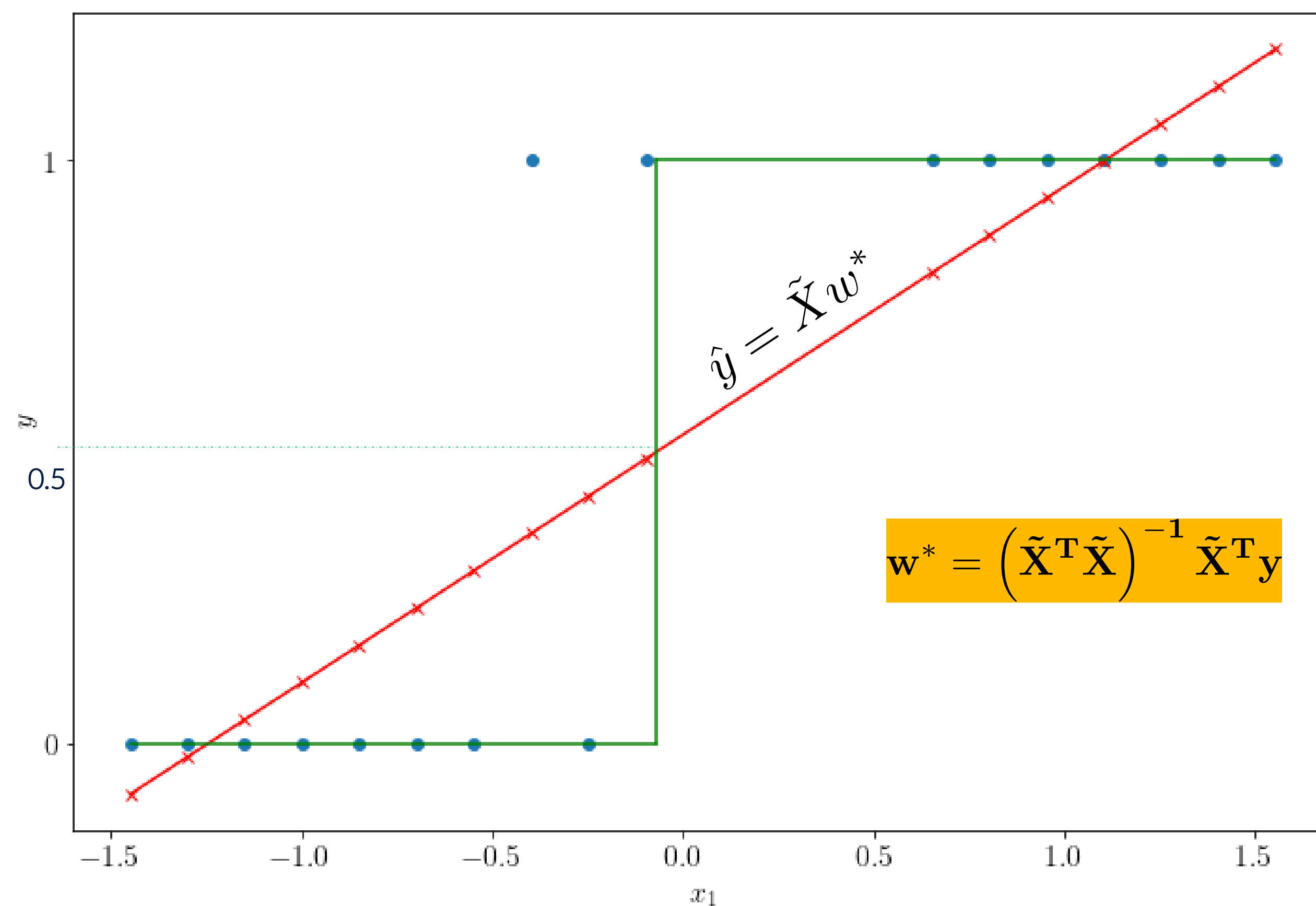
How about applying the linear regression framework seen earlier to fit such data?

We minimize the following cost function $J(w) = \frac{1}{N} \sum_{i=1}^N (\tilde{X}^i w - y^i)^2$



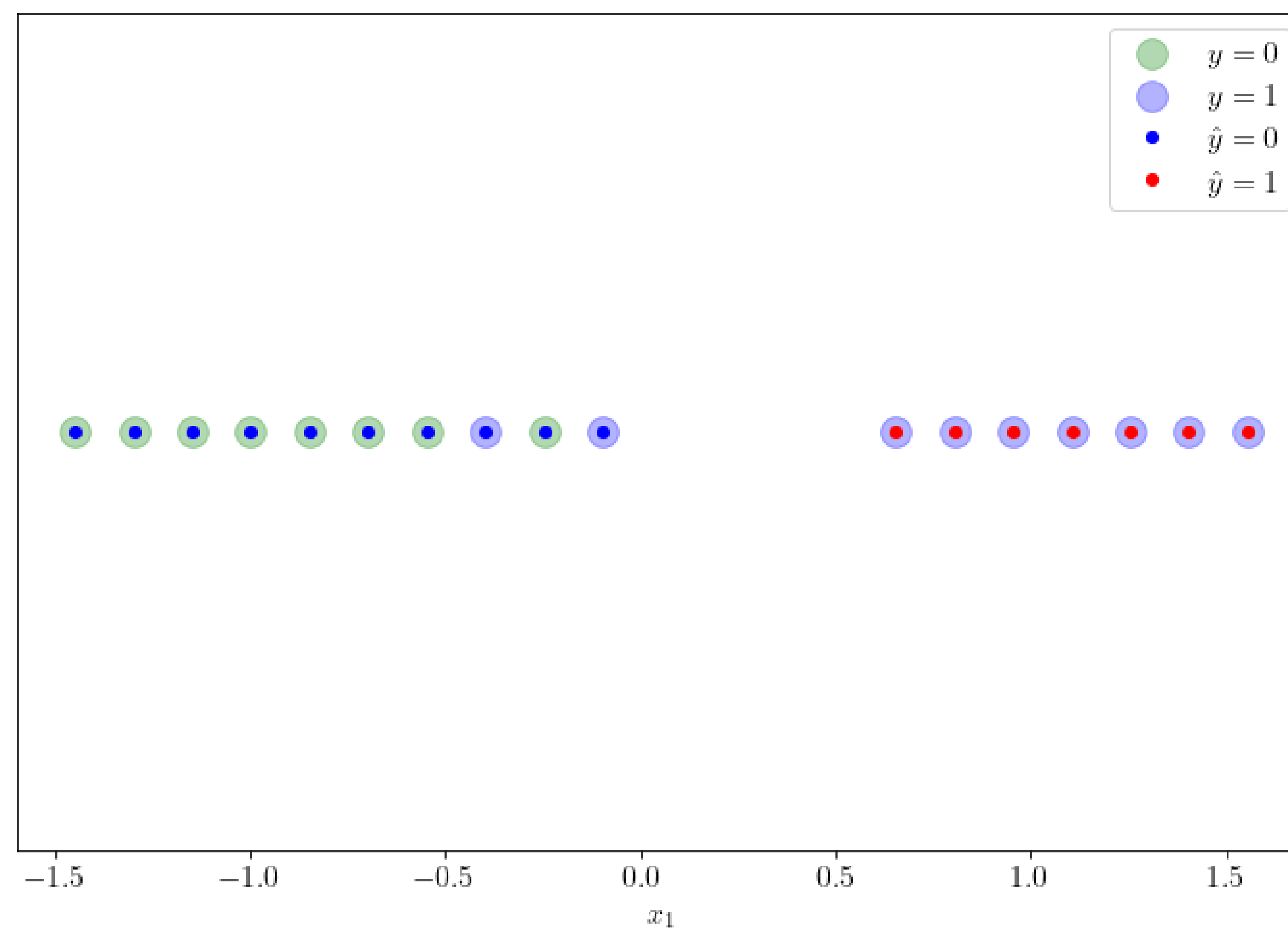
Proposition #1

And we apply a step function $\text{step}(\tilde{X}w^* - 0.5)$



Proposition #1

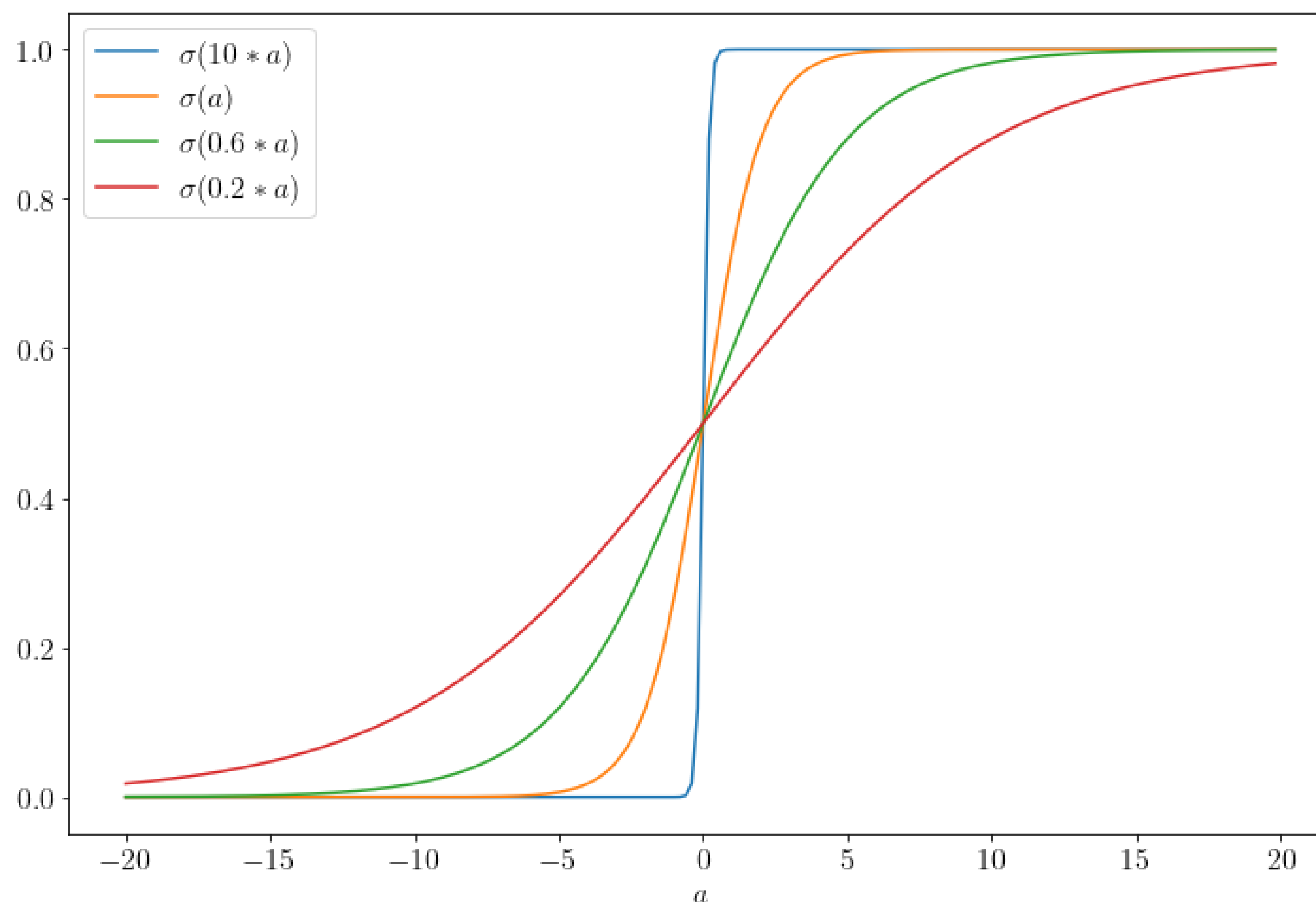
Results from the other perspective: 2 classification errors. In this case, 2 **false negatives**.



Proposition #2 – Intuition

How about integrating the step function in the cost function to minimize?

Or, better yet, a smooth approximation of the step function, the **logistic sigmoid** $\sigma(x) = \frac{1}{1 + e^{-x}}$



Proposition #2 – Formally

We'll explicitly model the probability $p := P(\hat{y} = 1|X)$

The probability takes values in $[0, 1]$. The range of a linear model is not bounded.

We'll extend the range to $[-\infty, +\infty]$, and approximate $\ln\left(\frac{p}{1-p}\right)$ by a linear model

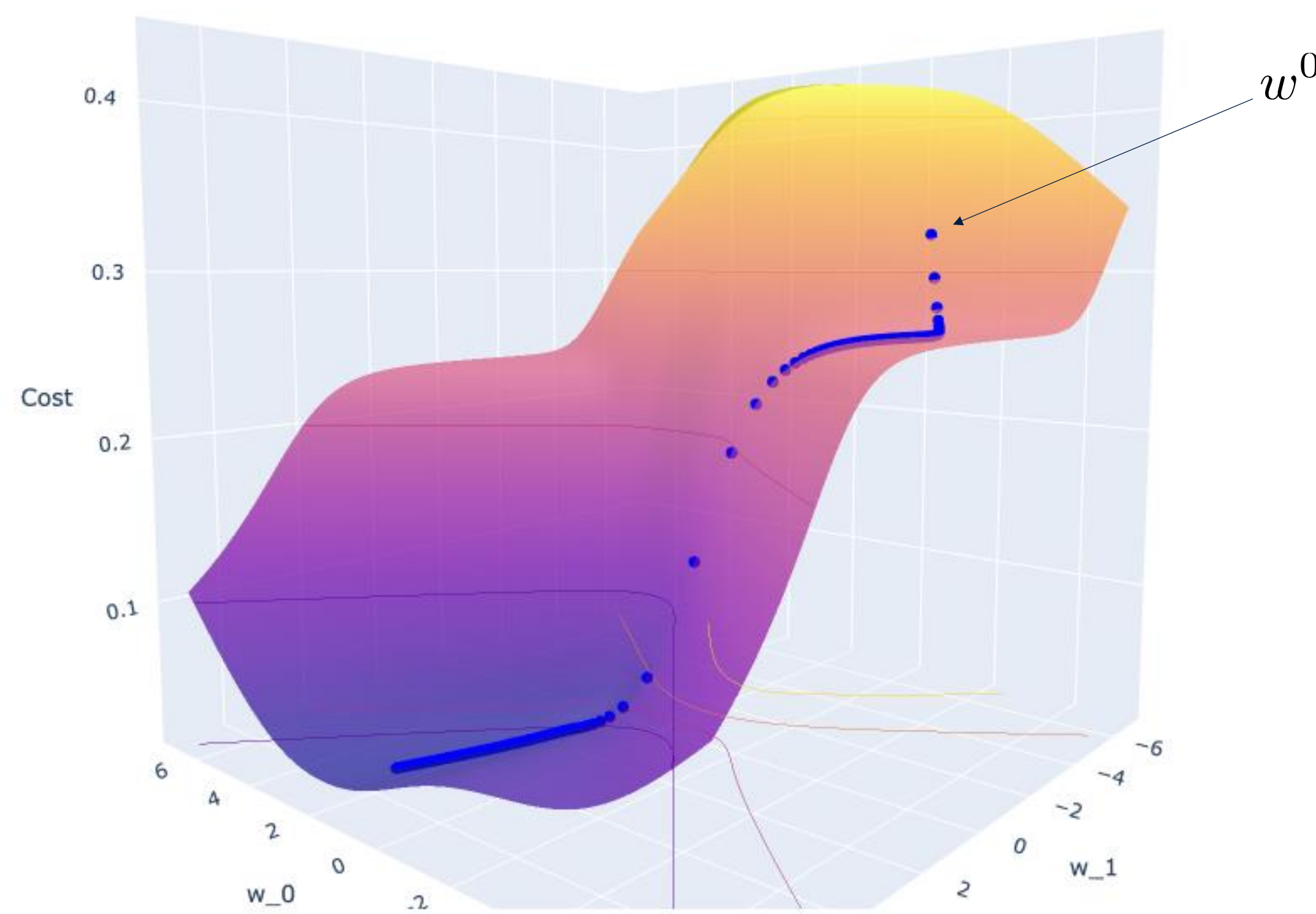
Hence, $\ln\left(\frac{p}{1-p}\right) = \tilde{X}w$, which can be written as:

$$p := P(\hat{y} = 1|X) = \frac{e^{\tilde{X}w}}{1 + e^{\tilde{X}w}} = \frac{1}{1 + e^{-\tilde{X}w}} = \sigma(\tilde{X}w)$$

Proposition #2

So we'll now minimize the following cost function $J_{\text{sl}}(w) = \frac{1}{2N} \sum_{i=1}^N \left(\sigma(\tilde{X}^i w) - y^i \right)^2$

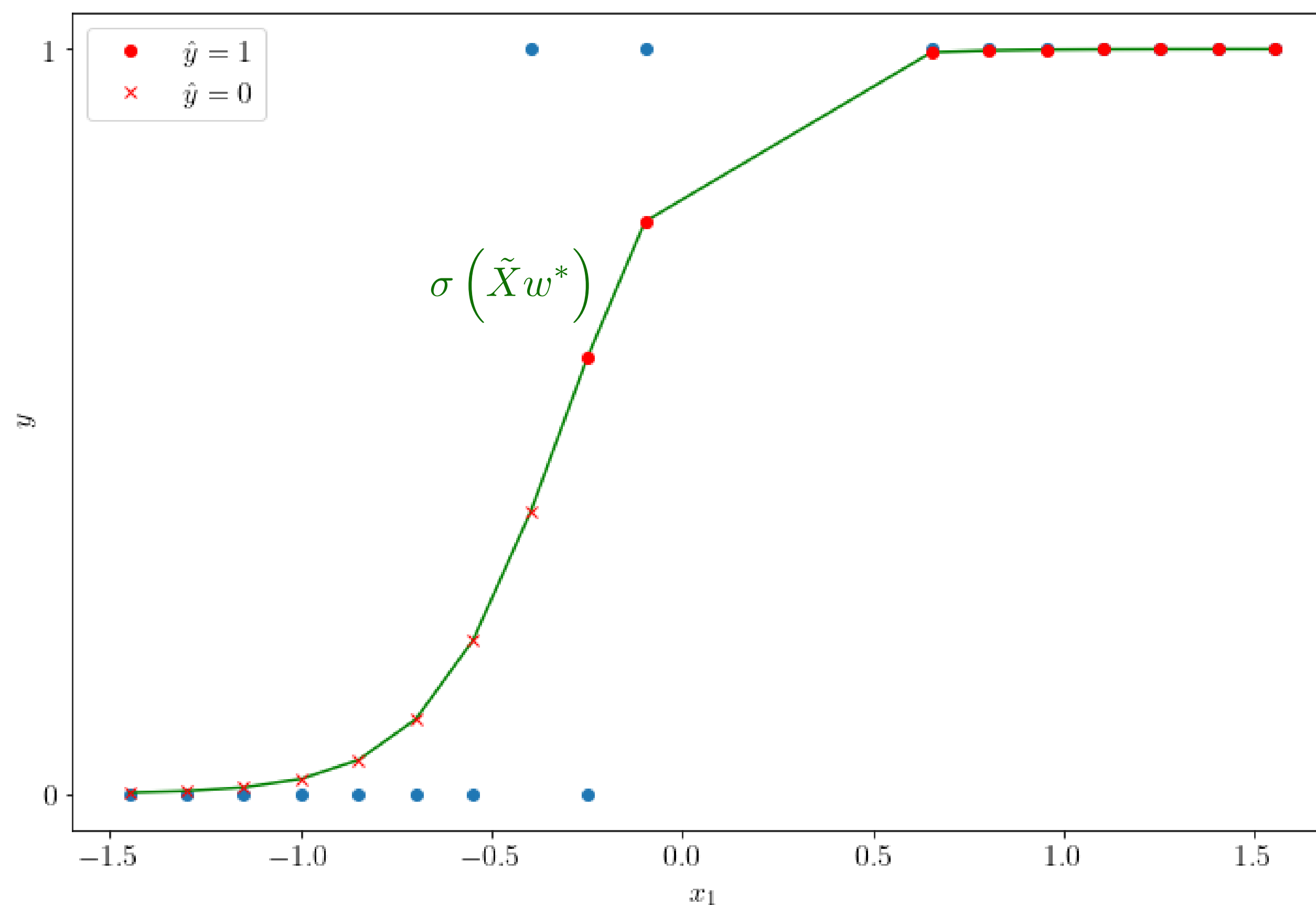
There is no closed-form solution for w^* . We'll use **gradient descent** (cf. notebook).



Proposition #2

So we'll now minimize the following cost function $J_{\text{sl}}(w) = \frac{1}{2N} \sum_{i=1}^N \left(\sigma(\tilde{X}^i w) - y^i \right)^2$

There is no closed-form solution for w^* . We'll use **gradient descent** (cf. notebook).



Proposition #3 – Cross entropy

A random variable follows a Bernoulli distribution if it only has two possible outcomes: 0 or 1.

$$P(Y = y \mid X) = p^y (1 - p)^{1-y}, \quad y = \{0, 1\}.$$

The likelihood of a Bernoulli distribution over N observations:

$$L(p) = \prod_{i=1}^N p^{y^i} (1 - p)^{1-y^i}$$

which we wish to maximize.

Recall that $p^i = P(\hat{y}^i = 1 \mid \tilde{X}^i) = \sigma(\tilde{X}^i w)$

We want to minimize the negative log-likelihood:

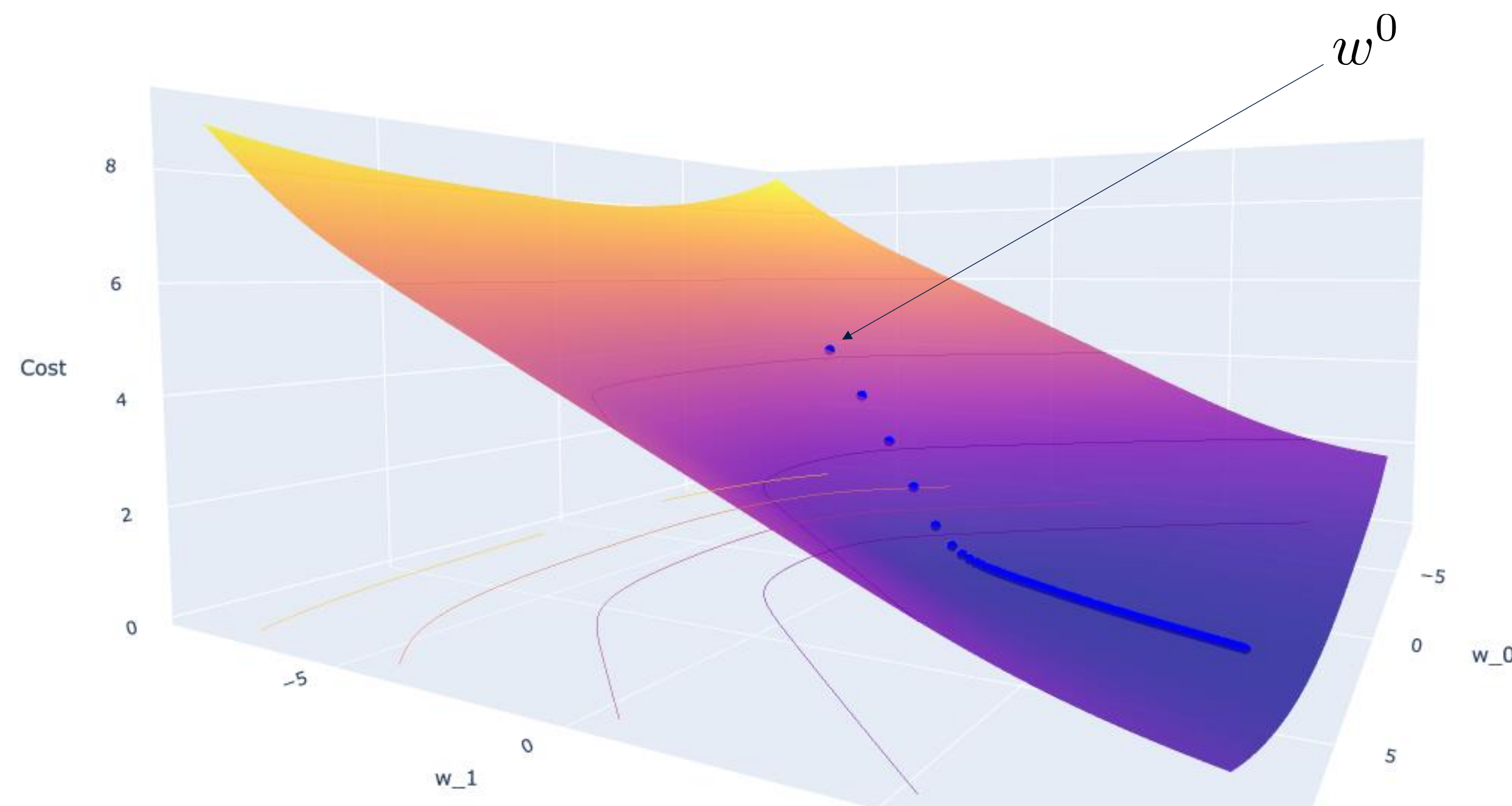
$$g_{\text{cel}} = -\frac{1}{N} \sum_{i=1}^N \left(y^i \ln(\sigma(\tilde{X}^i w)) + (1 - y^i) \ln(1 - \sigma(\tilde{X}^i w)) \right)$$

Proposition #3

We want to minimize the **negative of the log-likelihood** (or **cross entropy**):

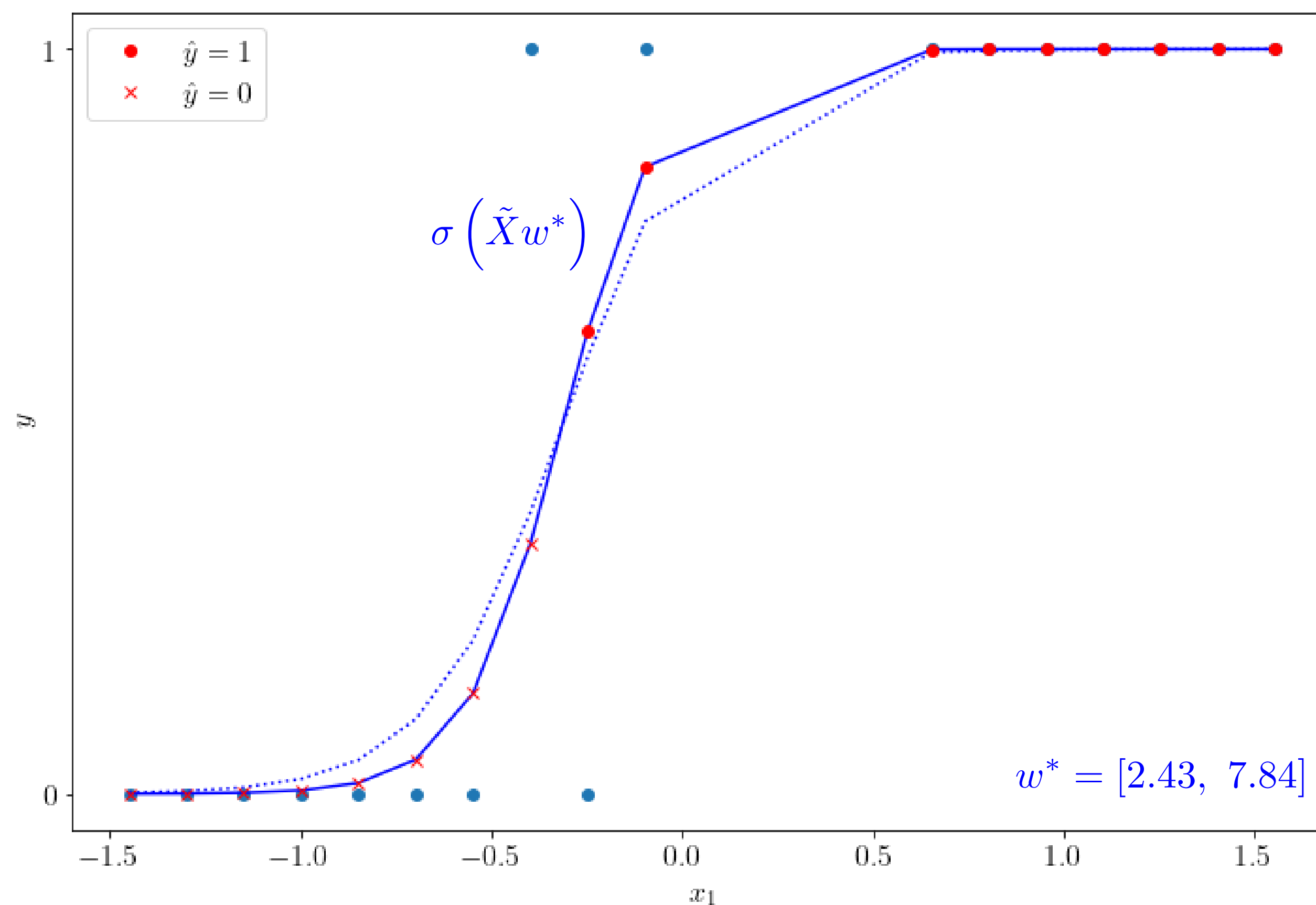
$$g_{\text{cel}} = -\frac{1}{N} \sum_{i=1}^N \left(y^i \ln \left(\sigma \left(\tilde{X}^i w \right) \right) + (1 - y^i) \ln \left(1 - \sigma \left(\tilde{X}^i w \right) \right) \right)$$

There is no closed-form solution for w^* . We'll use **gradient descent** (cf. notebook).



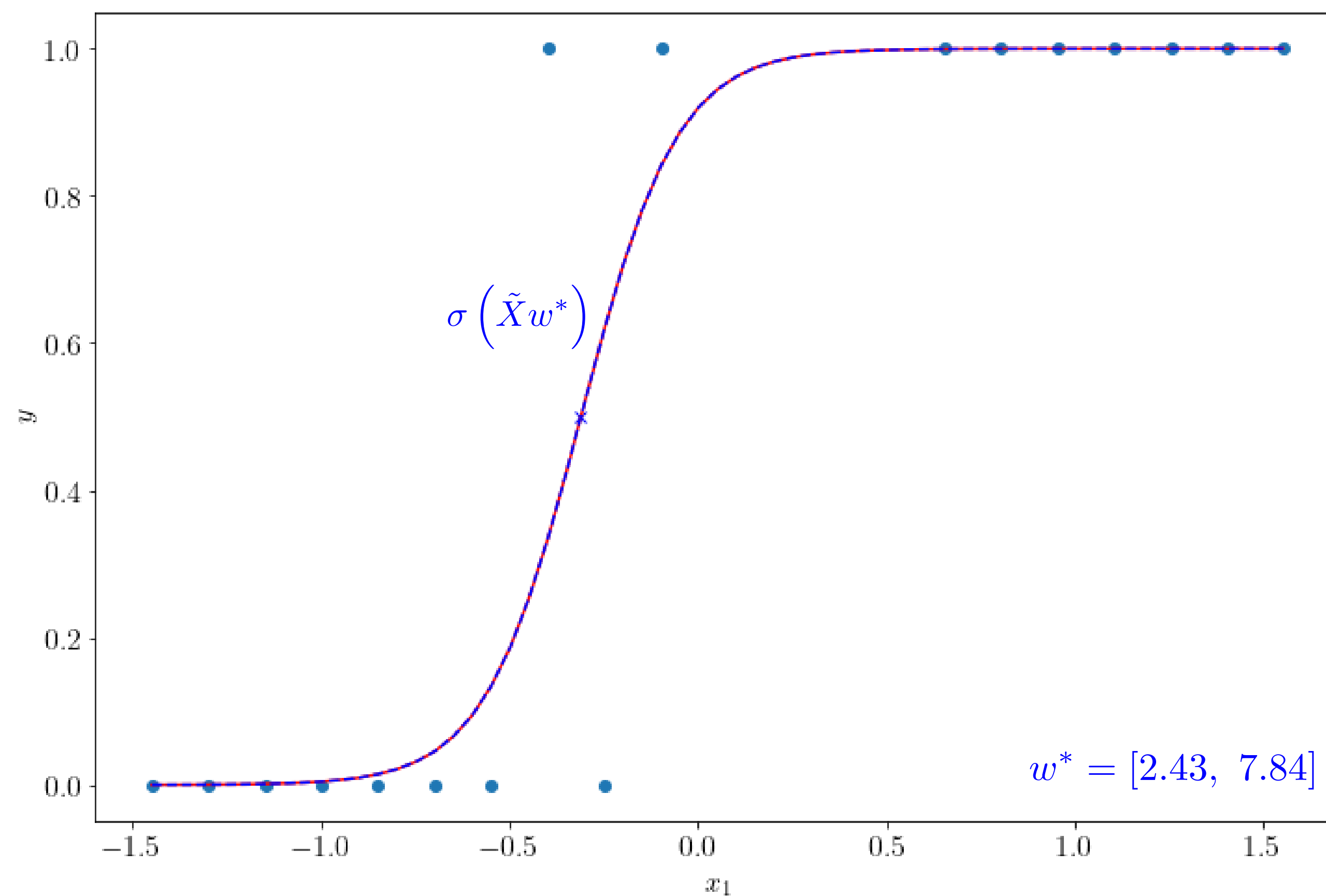
Minimizing the Cross Entropy

$$g_{\text{cel}} = -\frac{1}{N} \sum_{i=1}^N \left(y^i \ln \left(\sigma \left(\tilde{X}^i w \right) \right) + (1 - y^i) \ln \left(1 - \sigma \left(\tilde{X}^i w \right) \right) \right)$$



Now using scikit-learn

```
# Régularisation: 'none' ou 'l2'  
clf = sklearn.linear_model.LogisticRegression(solver='lbfgs', penalty='none')  
clf.fit(X, y.reshape(N,))  
  
w_0, w_1 = clf.intercept_[0], clf.coef_[0][0]
```





Performance evaluation Classification

Confusion matrix and statistics

		Predicted Class		
		Positive	Negative	
Actual Class	Positive	True Positive (TP)	False Negative (FN) Type II Error	Sensitivity $\frac{TP}{(TP + FN)}$ <i>Sensitivity: True Positive Rate or Recall</i>
	Negative	False Positive (FP) Type I Error	True Negative (TN)	Specificity $\frac{TN}{(TN + FP)}$ <i>(1 - Specificity): False Positive Rate</i>
		Precision $\frac{TP}{(TP + FP)}$	Negative Predictive Value $\frac{TN}{(TN + FN)}$	Accuracy $\frac{TP + TN}{(TP + TN + FP + FN)}$

$$\text{F1 Score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

Source: [Data Science and machine learning: The confusion matrix](#)

Exercise

- Calculate precision, sensitivity and specificity
- Calculate accuracy and F1 score

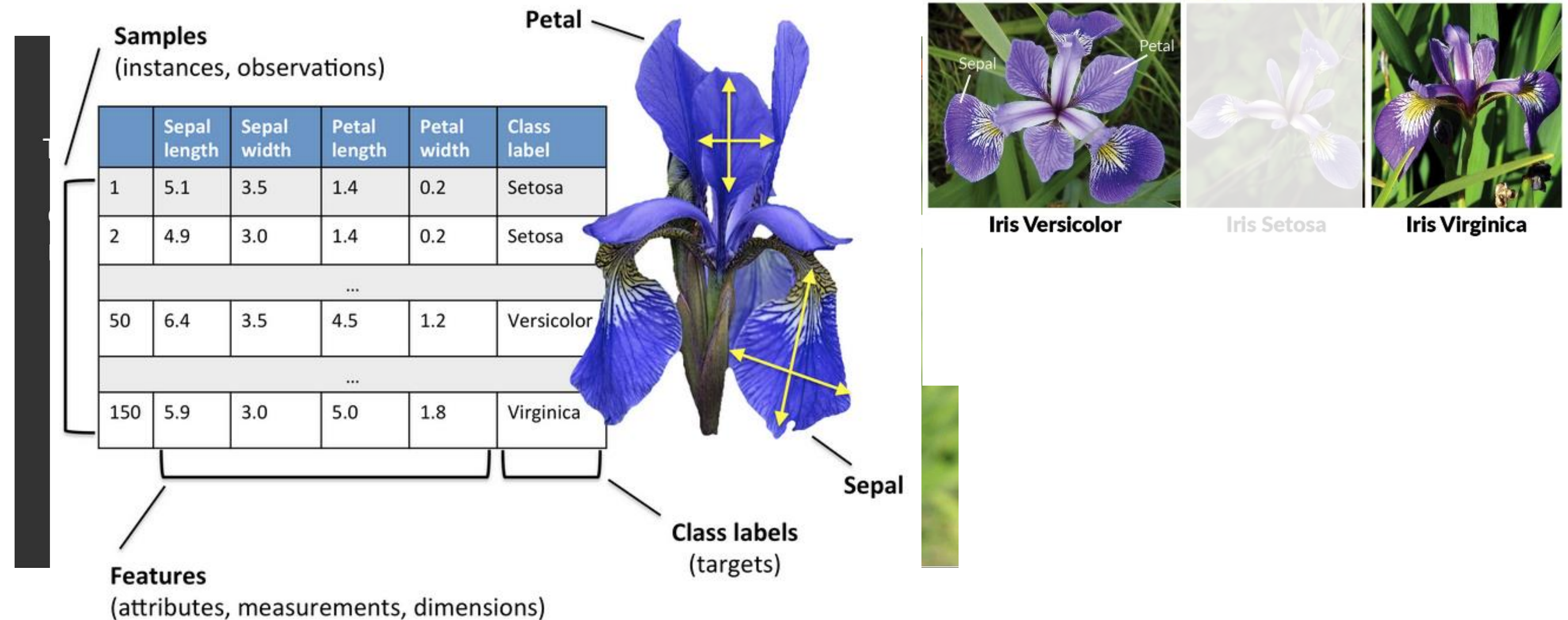
		Predicted Class	
		Spam	Non-Spam
Actual Class	Spam	TP=45	FN=20
	Non-Spam	FP=5	TN=30

Source: [Data Science and machine learning: The confusion matrix](#)

Real example using scikit-learn

27

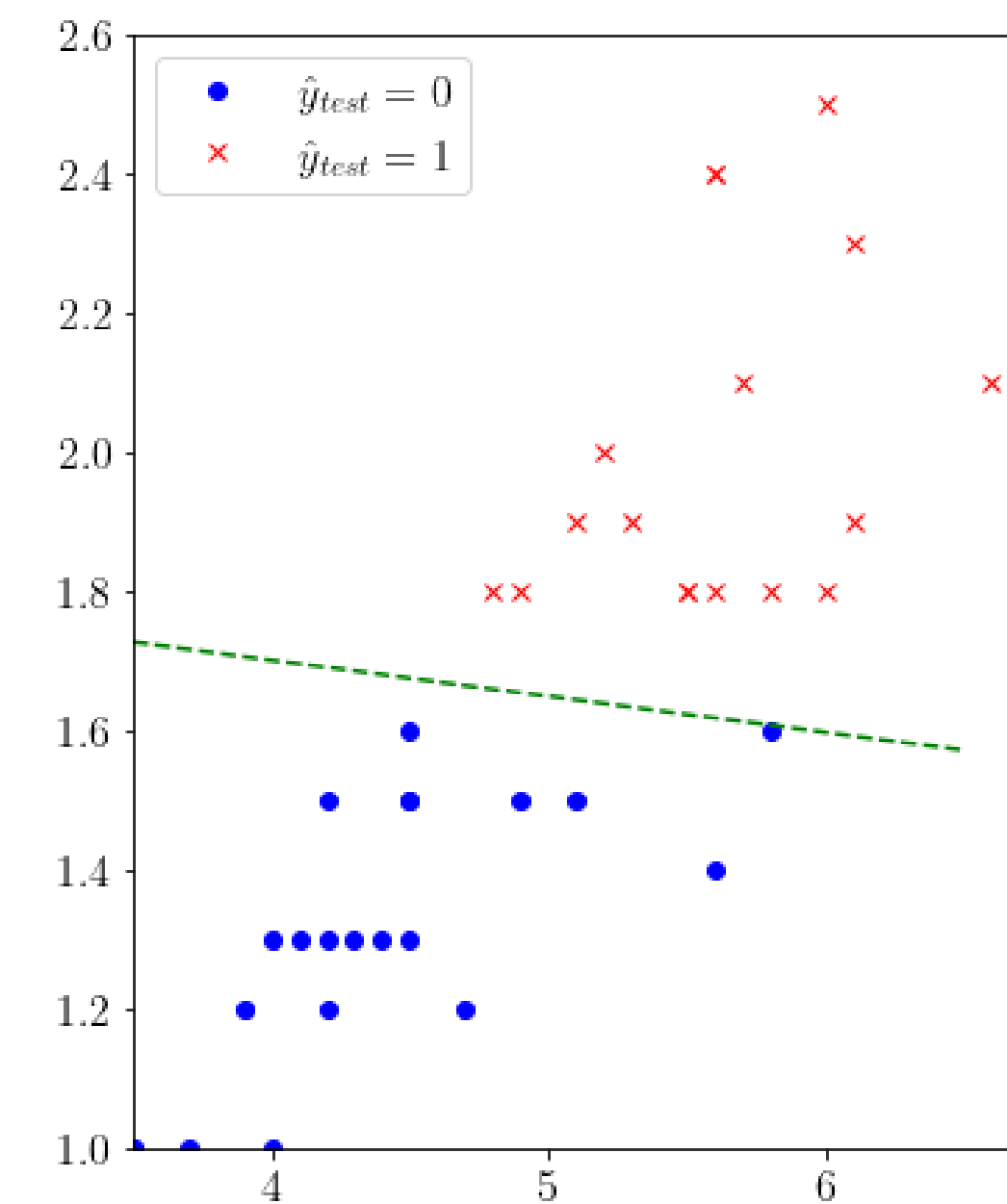
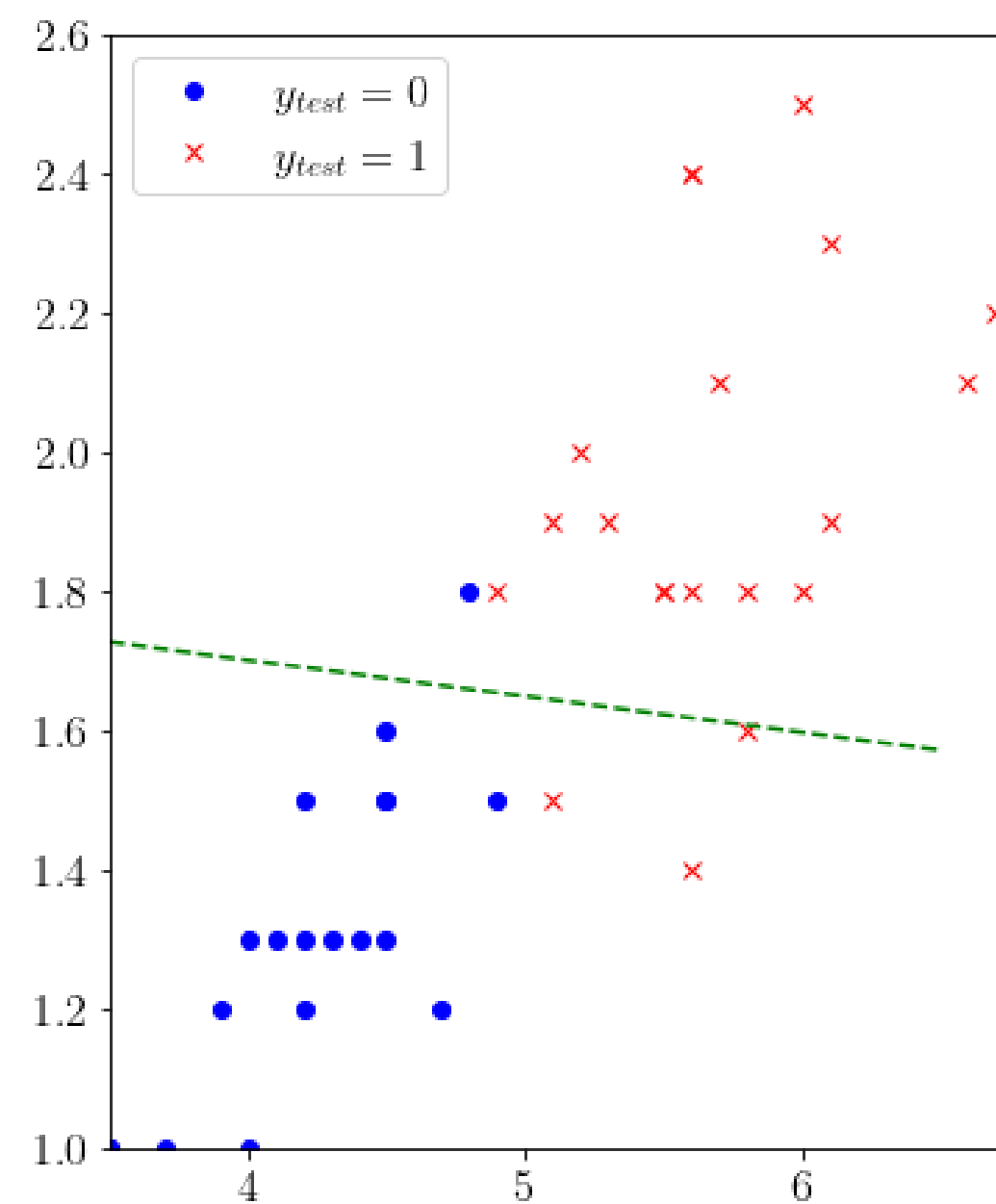
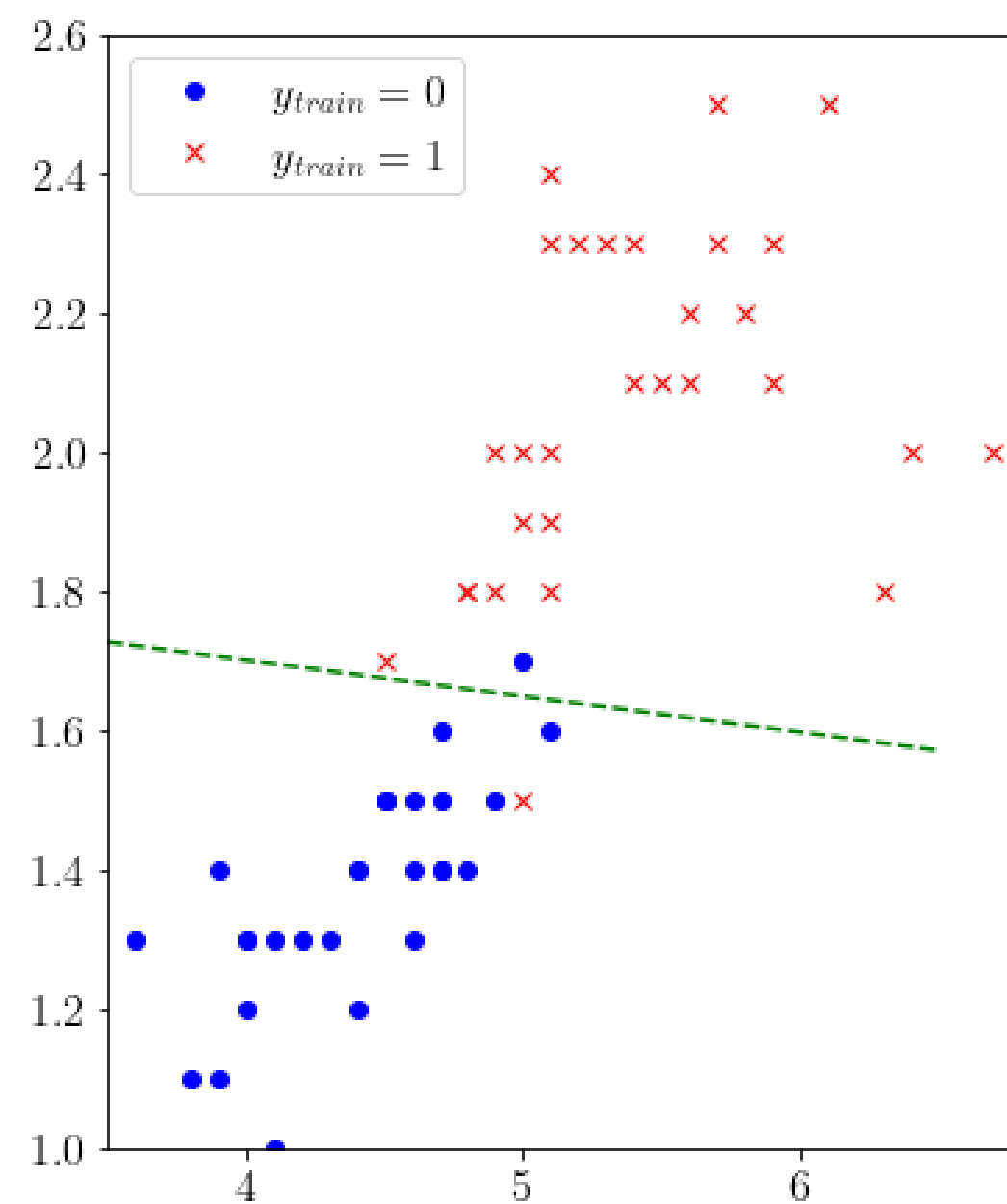
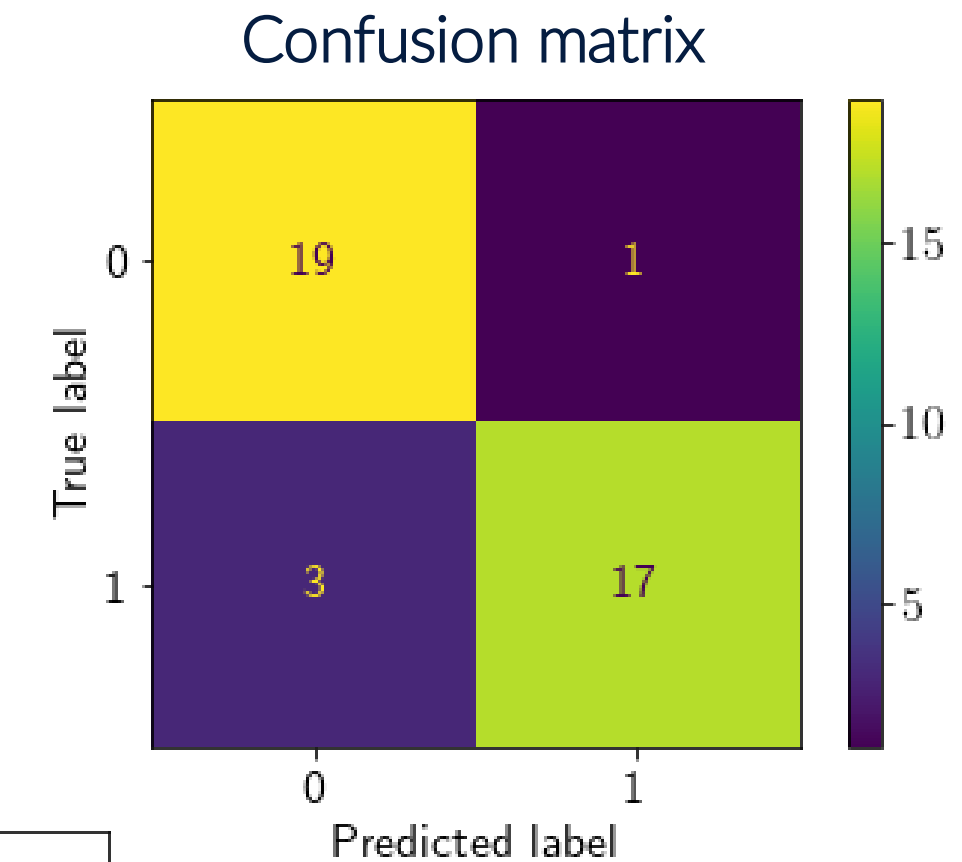
UCI Machine Learning Repository



Iris dataset {versicolor: 0, virginica: 1}

28

- Using train-test-split, 30% test
- Using logistic regression (no regularization)



What did I learn?

- Definition of (binary) classification (vs regression)
- **Logistic regression** using various cost functions
- In particular, the **cross entropy cost** (or negative log-likelihood)
- **Confusion matrix** for classification performance

QUESTIONS ?

