

# **Maximum Likelihood Estimation**

## **and**

# **Generalized Linear Models**

**Prepared by: Joseph Bakarji**

# Probabilistic Interpretation of Linear Regression

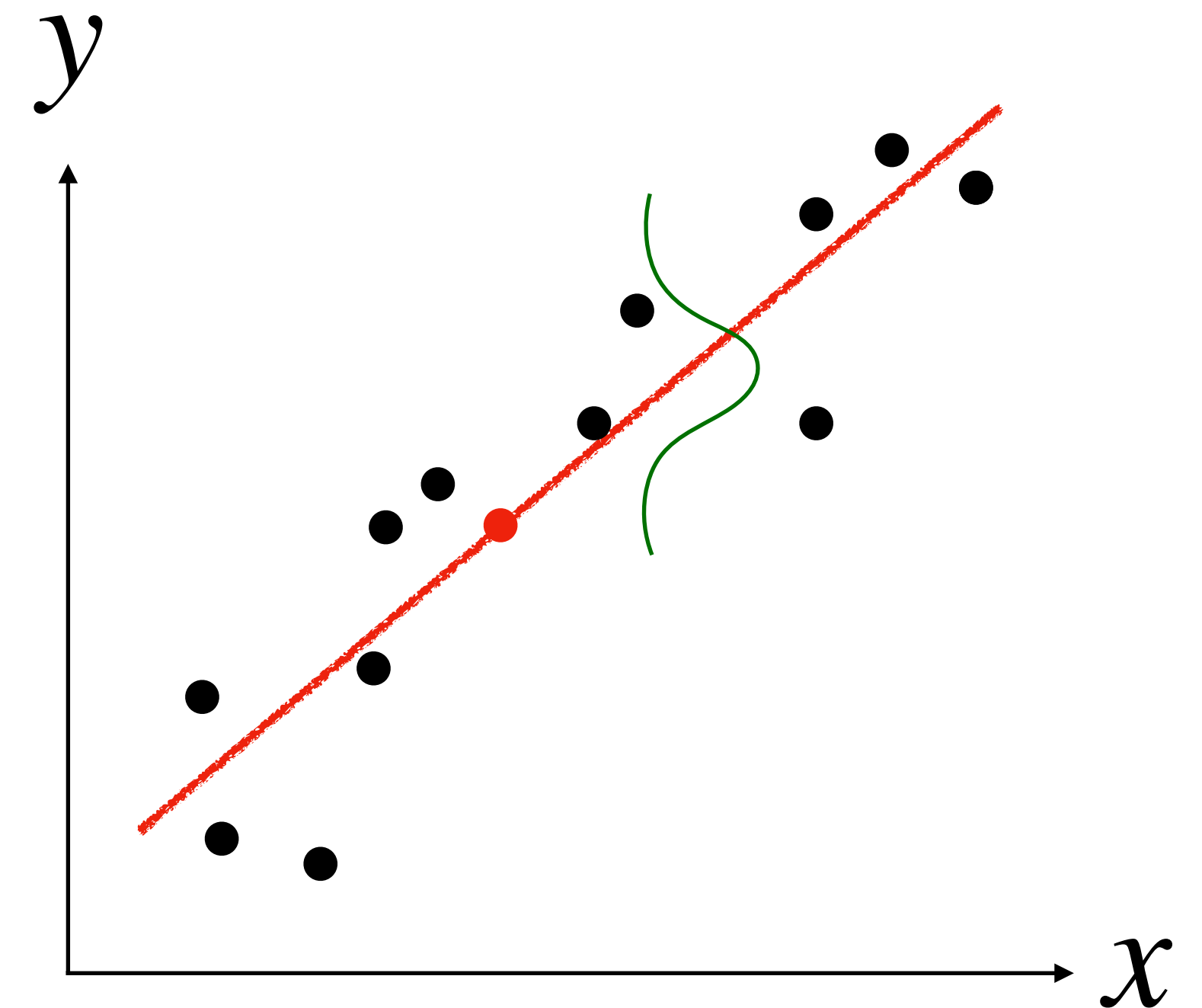
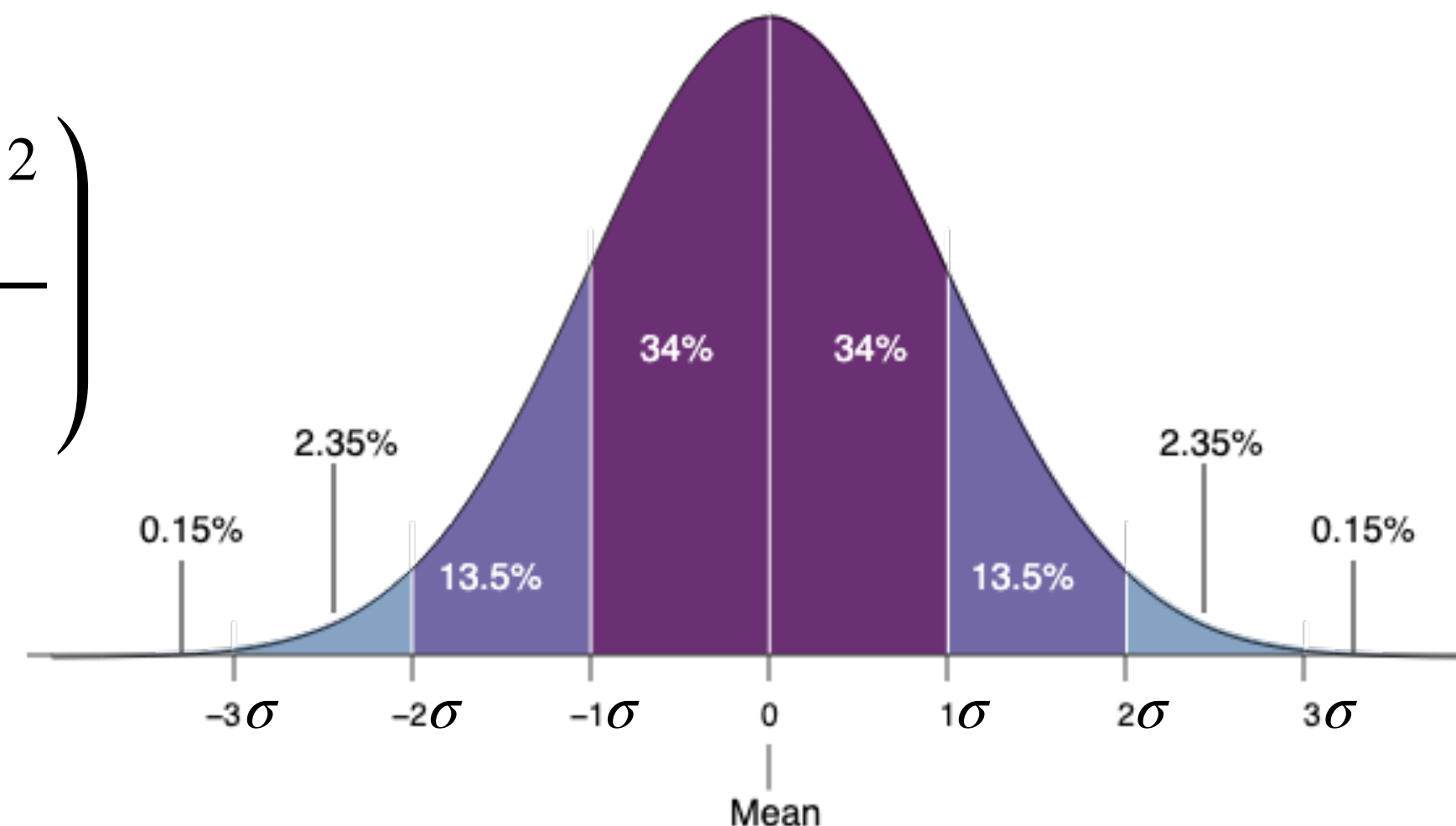
Assume **noise** is normally distributed around model

$$y^{(i)} = \theta^\top x^{(i)} + \varepsilon^{(i)}$$

Normally distributed

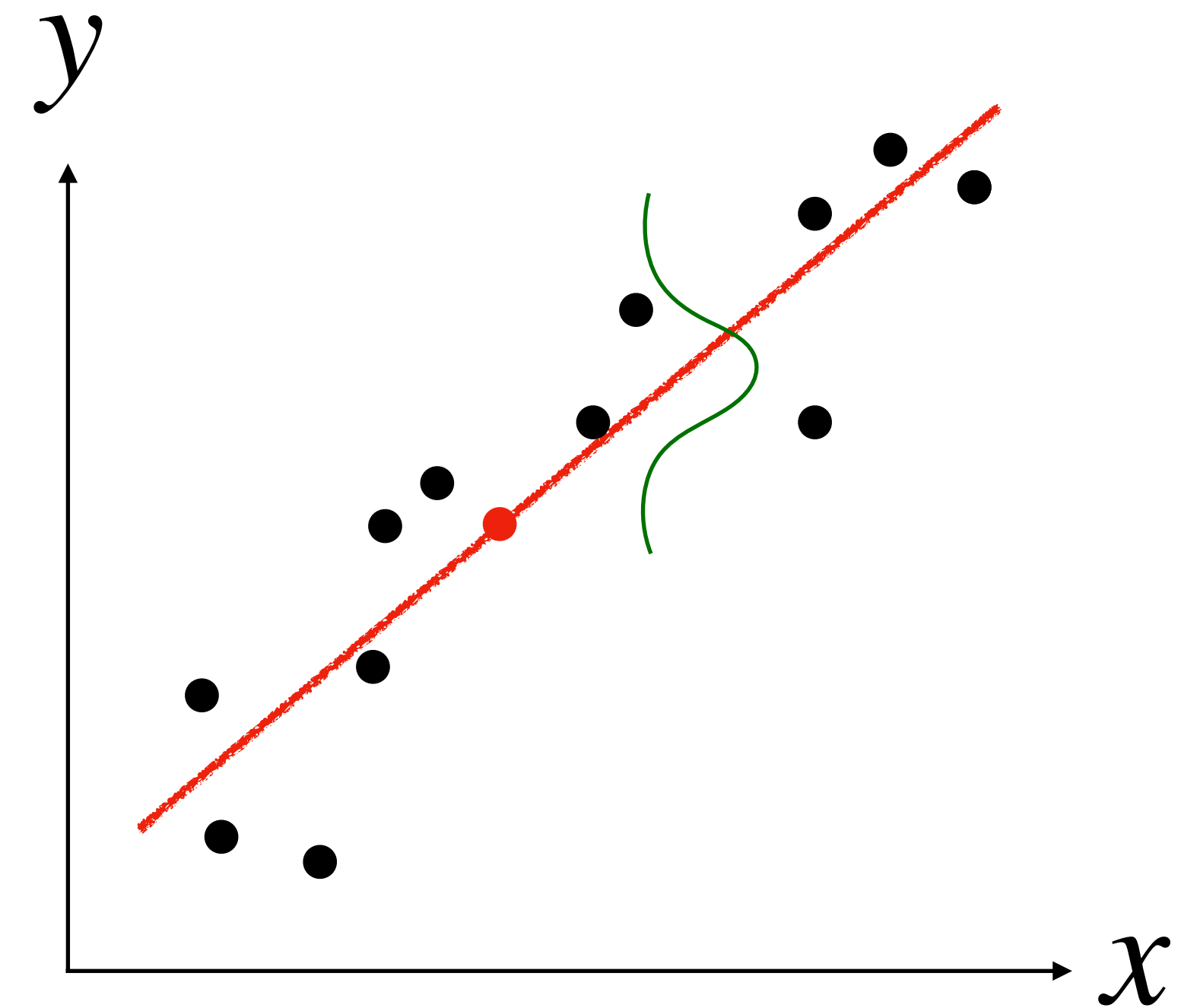
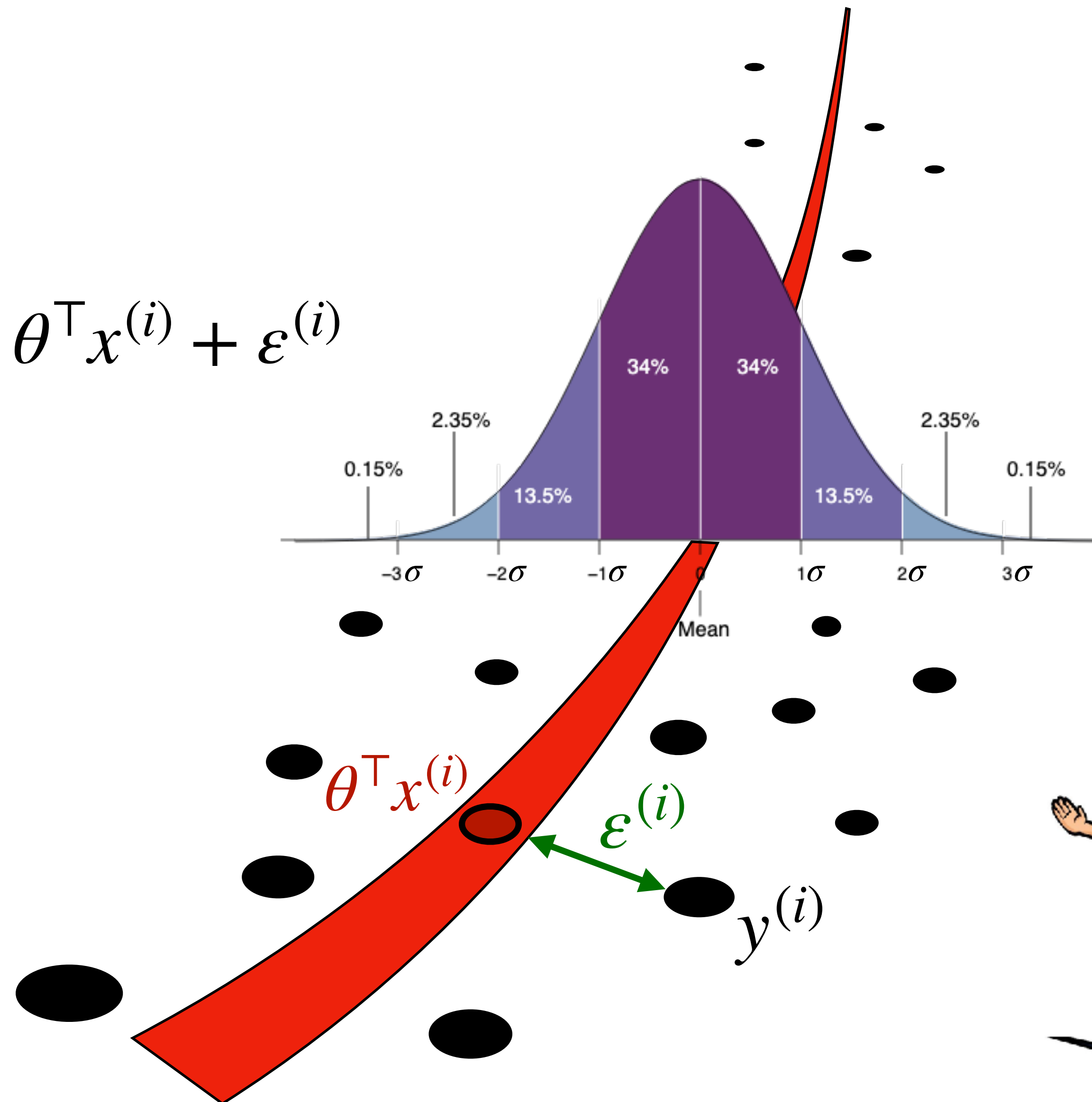
$$\mathcal{N}(0, \sigma^2)$$

$$p(\varepsilon^{(i)}) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(\varepsilon^{(i)})^2}{2\sigma^2}\right)$$



# Probabilistic Interpretation of Linear Regression

$$y^{(i)} = \theta^T x^{(i)} + \varepsilon^{(i)}$$



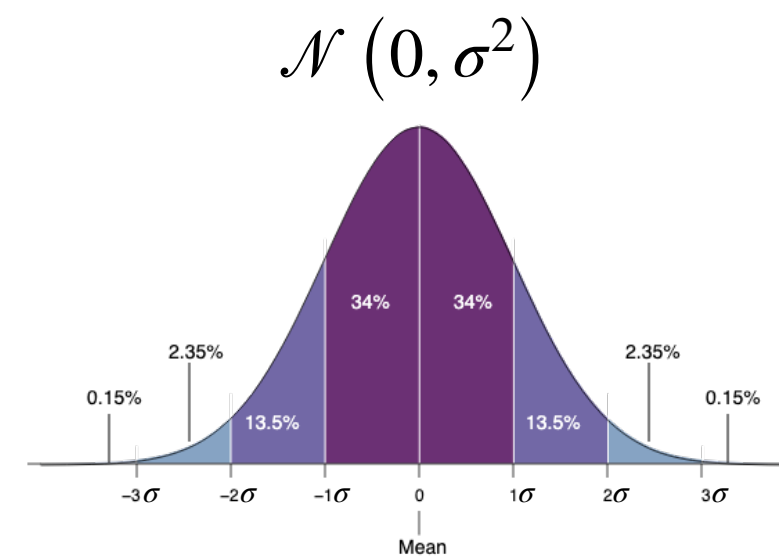


# Probabilistic Interpretation of Linear Regression

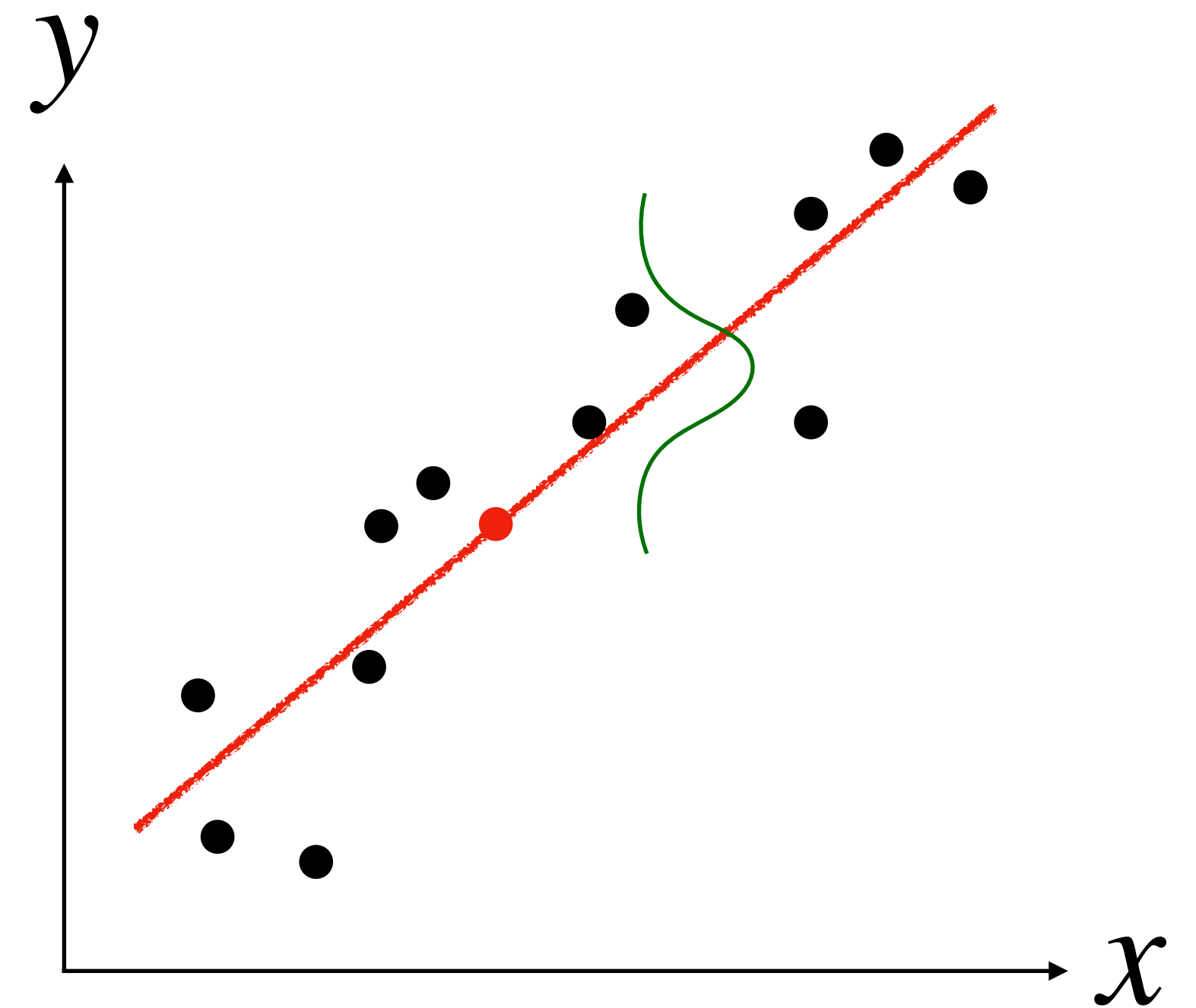
Assume noise is normally distributed around model

$$y^{(i)} = \theta^\top x^{(i)} + \epsilon^{(i)}$$

$$p(\epsilon^{(i)}) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(\epsilon^{(i)})^2}{2\sigma^2}\right)$$



$$p(y^{(i)} | x^{(i)}; \theta) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(y^{(i)} - \theta^\top x^{(i)})^2}{2\sigma^2}\right)$$

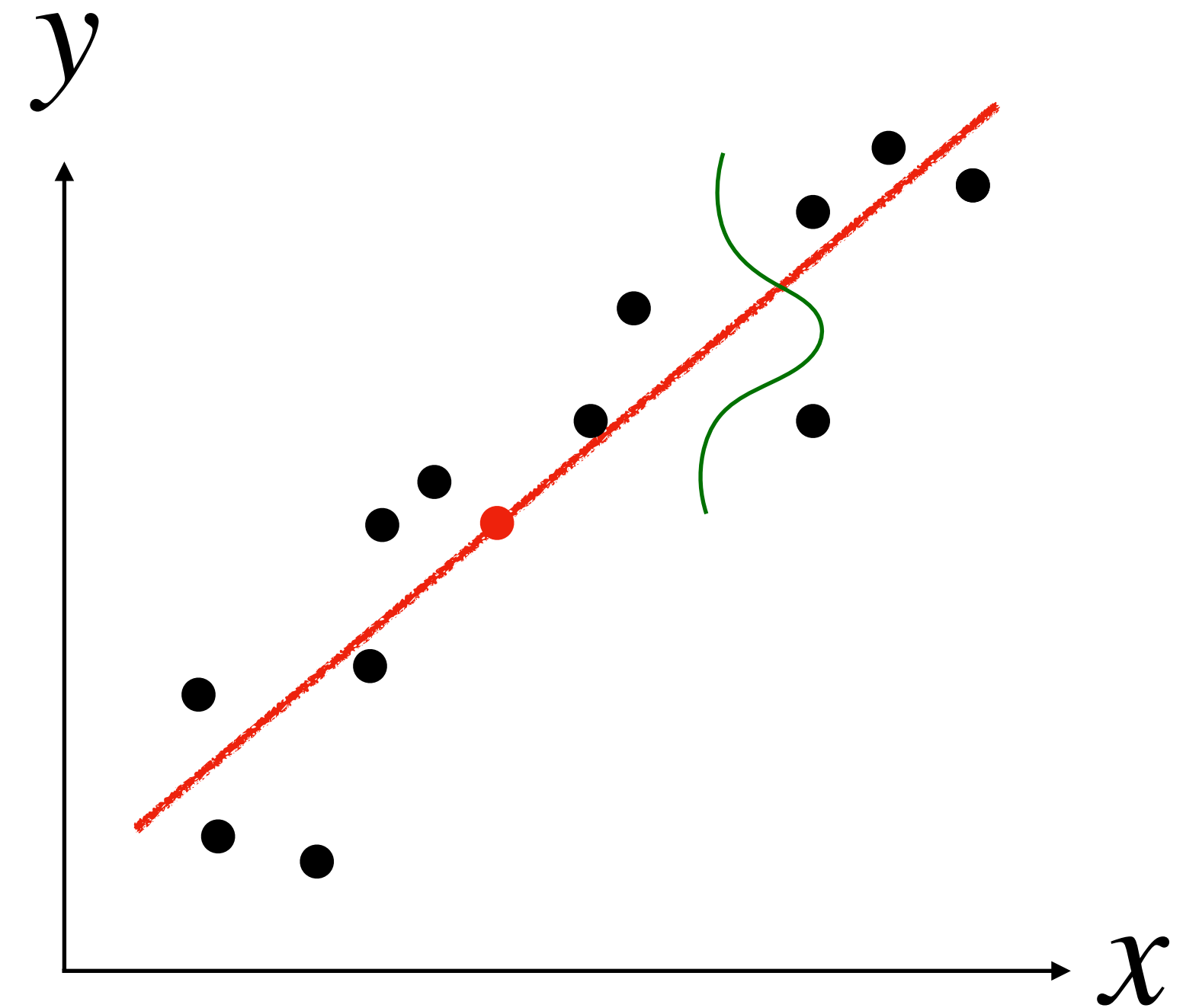


# Likelihood of output given input

$$L(\theta) = \prod_{i=1}^n p(y^{(i)} | x^{(i)}; \theta) \quad \text{Independent and Identically Distributed (IID)}$$
$$= \prod_{i=1}^n \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(y^{(i)} - \theta^\top x^{(i)})^2}{2\sigma^2}\right)$$

## Log-likelihood

$$\mathcal{L}(\theta) = \log L(\theta)$$
$$= \log \prod_{i=1}^n \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(y^{(i)} - \theta^\top x^{(i)})^2}{2\sigma^2}\right)$$
$$= \sum_{i=1}^n \log \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(y^{(i)} - \theta^\top x^{(i)})^2}{2\sigma^2}\right) = n \log \frac{1}{\sqrt{2\pi}\sigma} - \frac{1}{2\sigma^2} \sum_{i=1}^n (y^{(i)} - \theta^\top x^{(i)})^2$$



# Maximize **Log-likelihood**

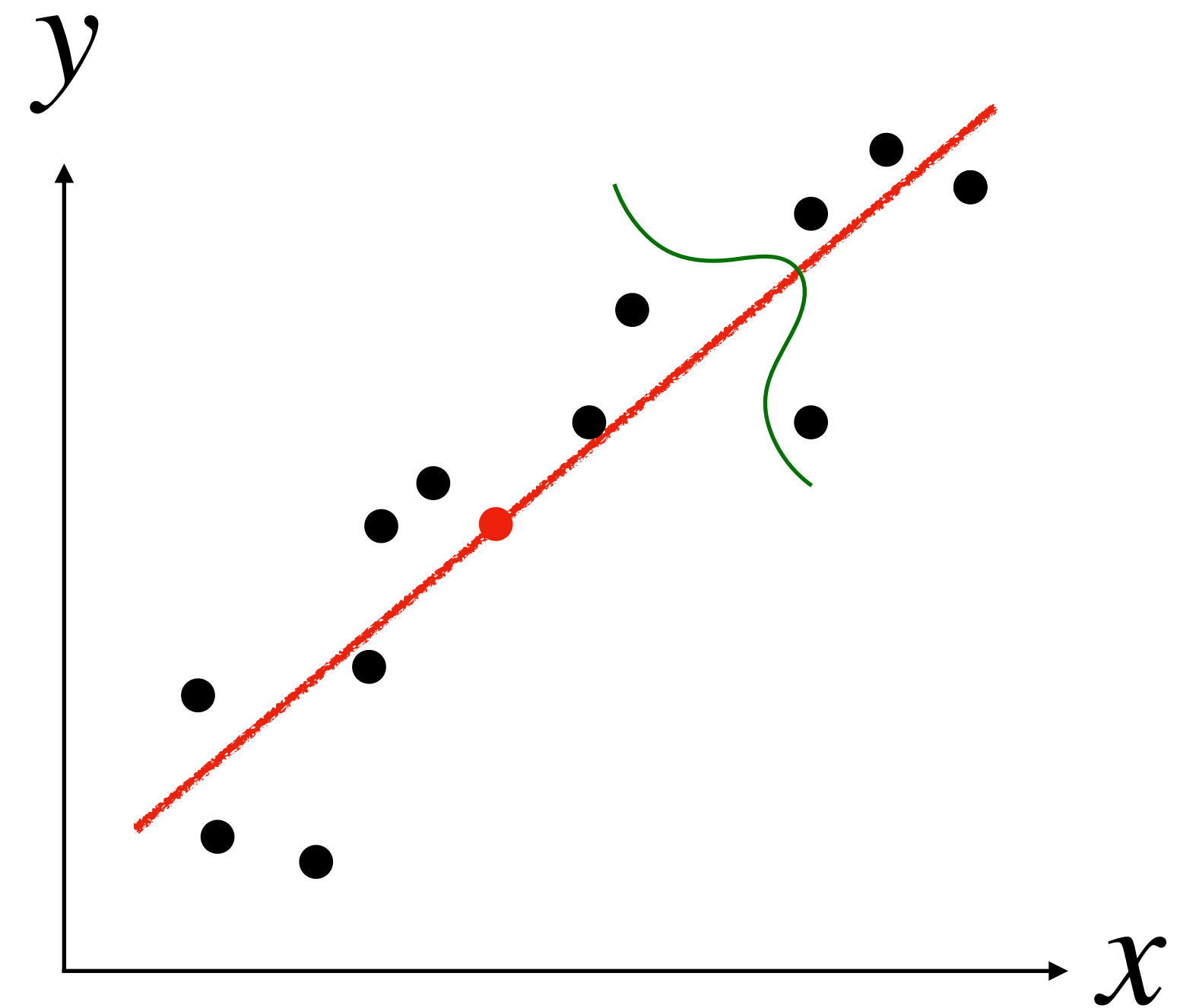
$$\mathcal{L}(\theta) = \log L(\theta)$$

$$= \sum_{i=1}^n \log \frac{1}{\sqrt{2\pi}\sigma} \exp \left( -\frac{(y^{(i)} - \theta^\top x^{(i)})^2}{2\sigma^2} \right)$$

$$= n \log \frac{1}{\sqrt{2\pi}\sigma} - \frac{1}{2\sigma^2} \sum_{i=1}^n (y^{(i)} - \theta^\top x^{(i)})^2$$

$$\text{Maximize } \mathcal{L}(\theta) \longrightarrow \text{Minimize } \frac{1}{2} \sum_{i=1}^n (y^{(i)} - \theta^\top x^{(i)})^2$$

**What if the noise is not Gaussian?**

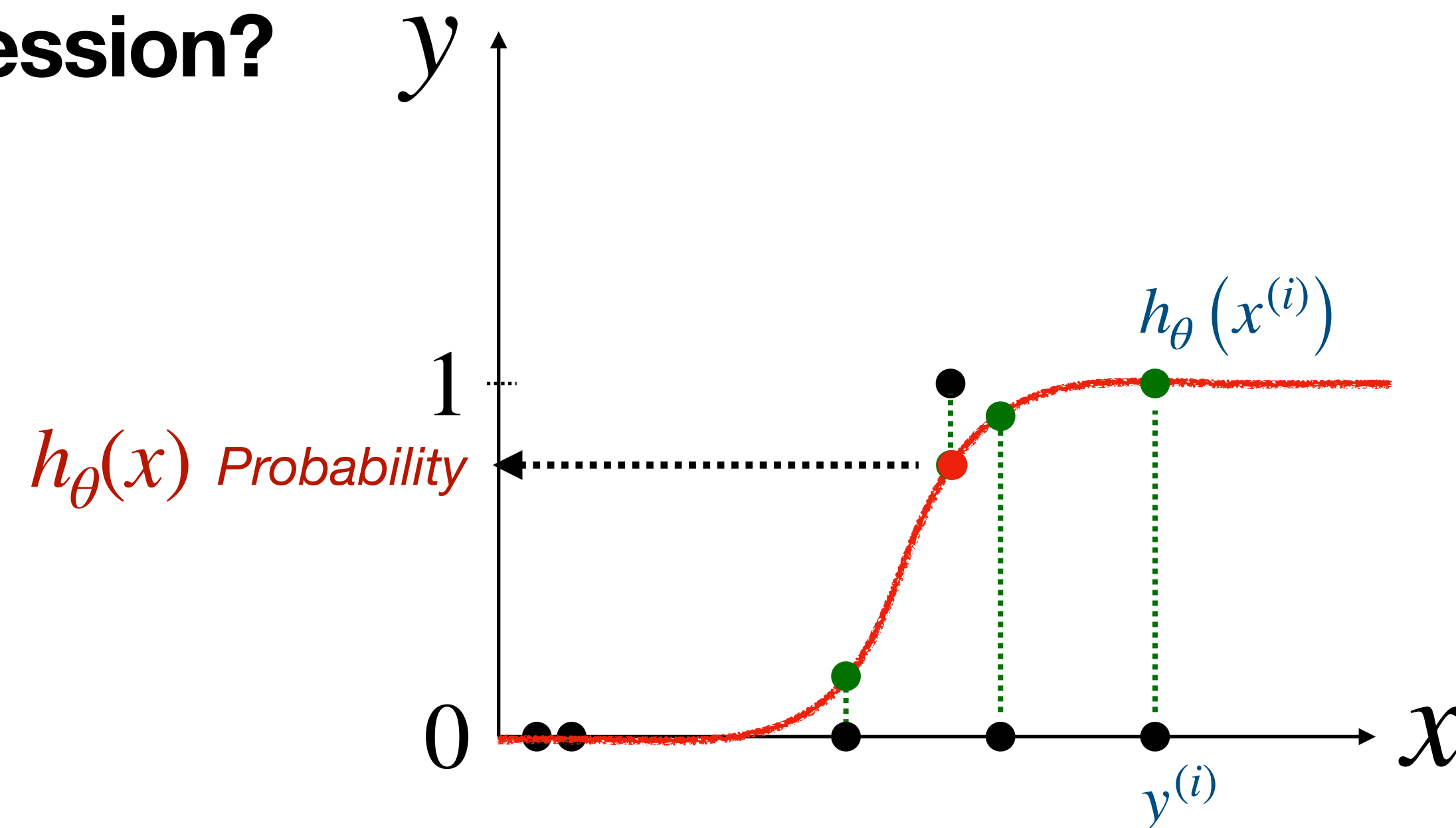




## Why not Least Squares in Logistic Regression?

$$y = h_{\theta}(x)$$

$$h_{\theta}(x) = \frac{1}{1 + e^{-(\theta^{\top} x)}} = \sigma(\theta^{\top} x)$$



## Probability of output given input

$$P(y = 1 \mid x; \theta) = h_{\theta}(x)$$

$$P(y = 0 \mid x; \theta) = 1 - h_{\theta}(x)$$



True label

$$p(y \mid x; \theta) = (h_{\theta}(x))^y (1 - h_{\theta}(x))^{1-y}$$

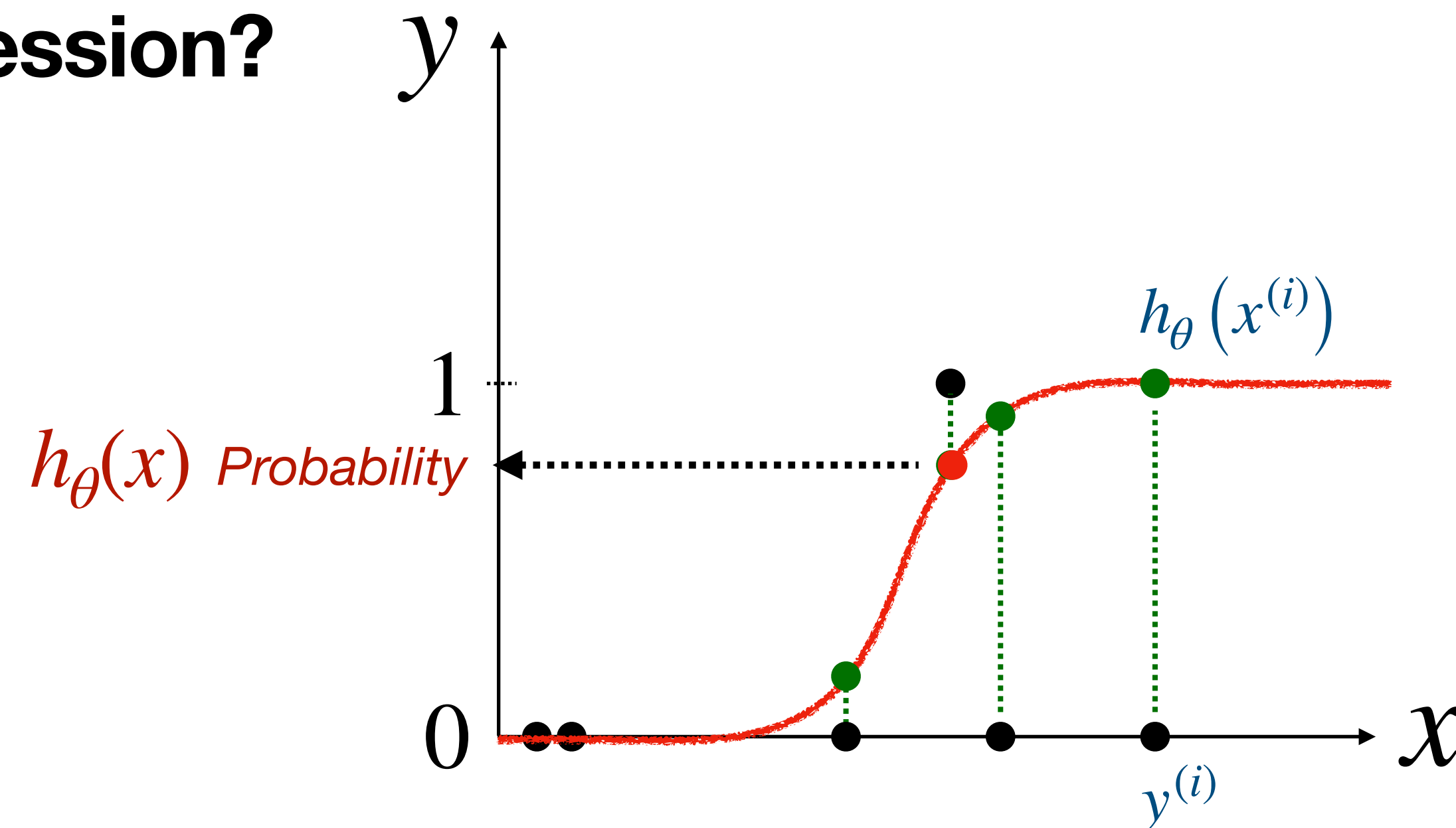
**Likelihood!**

**For Bernoulli Distributed Noise**

## Why not Least Squares in Logistic Regression?

$$y = h_{\theta}(x)$$

$$h_{\theta}(x) = \frac{1}{1 + e^{-(\theta^{\top} x)}} = \sigma(\theta^{\top} x)$$



## Probability of output given input

$$P(y = 1 \mid x; \theta) = \sigma(\theta^{\top} x)$$

$$P(y = 0 \mid x; \theta) = 1 - \sigma(\theta^{\top} x)$$



True label

$$p(y \mid x; \theta) = (\sigma(\theta^{\top} x))^y (1 - \sigma(\theta^{\top} x))^{1-y}$$

**Likelihood!**

**For Bernoulli Distributed Noise**



# Bernoulli Distribution

## Properties [\[edit\]](#)

---

If  $X$  is a random variable with a Bernoulli distribution, then:

$$\Pr(X = 1) = p = 1 - \Pr(X = 0) = 1 - q.$$

The [probability mass function](#)  $f$  of this distribution, over possible outcomes  $k$ , is

$$f(k; p) = \begin{cases} p & \text{if } k = 1, \text{ [3]} \\ q = 1 - p & \text{if } k = 0. \end{cases}$$

This can also be expressed as

$$f(k; p) = p^k (1 - p)^{1-k} \quad \text{for } k \in \{0, 1\}$$

or as

$$f(k; p) = pk + (1 - p)(1 - k) \quad \text{for } k \in \{0, 1\}.$$

The Bernoulli distribution is a special case of the [binomial distribution](#) with  $n = 1$ .[\[4\]](#)

# Define **Log-likelihood**

## Likelihood

$$p(y | x; \theta) = \left( h_{\theta}(x) \right)^y \left( 1 - h_{\theta}(x) \right)^{1-y} \quad \text{for all } (x, y) \text{ pair}$$

$$\begin{aligned} L(\theta) &= \prod_{i=1}^n p \left( y^{(i)} | x^{(i)}; \theta \right) \\ &= \prod_{i=1}^n h_{\theta} \left( x^{(i)} \right)^{y^{(i)}} \left( 1 - h_{\theta} \left( x^{(i)} \right) \right)^{1-y^{(i)}} \end{aligned}$$

log  
↓

$$\mathcal{L}(\theta) = \sum_{i=1}^n y^{(i)} \log h_{\theta} \left( x^{(i)} \right) + \left( 1 - y^{(i)} \right) \log \left( 1 - h_{\theta} \left( x^{(i)} \right) \right)$$

# Maximize **Log-likelihood**

$$\mathcal{L}(\theta) = \sum_{i=1}^n y^{(i)} \log h_{\theta}(x^{(i)}) + (1 - y^{(i)}) \log (1 - h_{\theta}(x^{(i)}))$$

## Update rule?

**while** not converged:

$$\theta_j := \theta_j + \alpha \frac{\partial \mathcal{L}(\theta)}{\partial \theta_j}$$

**Derive**

## Gradient Descent

**for**  $t = 1 \dots T$ :

**for** all parameters  $j$ :

$$\theta_j := \theta_j - \alpha \sum_{i=1}^n \left( h_{\theta}(x^{(i)}) - y^{(i)} \right) x_j^{(i)}$$



# Maximize **Log-likelihood**

$$\mathcal{L}(\theta) = \sum_{i=1}^n y^{(i)} \log h_{\theta}(x^{(i)}) + (1 - y^{(i)}) \log (1 - h_{\theta}(x^{(i)}))$$

## Update rule

**while** not converged:

$$\theta := \theta + \alpha \nabla_{\theta} \mathcal{L}(\theta)$$



**Derive**

## Gradient Descent

**for** t = 1...T:

$$\theta := \theta - \alpha \sum_{i=1}^n \left( h_{\theta}(x^{(i)}) - y^{(i)} \right) x^{(i)}$$



Same as **linear regression**



# Generalized Linear Models

**Gaussian Distribution**



**Linear Regression**

**Bernoulli Distribution**



**Logistic Regression**

**Update rule**

$$\theta := \theta - \alpha \sum_{i=1}^n \left( h_{\theta} (x^{(i)}) - y^{(i)} \right) x^{(i)}$$

**Are there other distributions that lead to the same update rule?**

# Exponential Family

Family of distributions for which we can derive **the same update rule**

**Assumption:**  $p(y | x; \theta)$  is an exponential family

Data ←

$$p(y; \eta) = b(y) \exp \{ \eta^\top y - a(\eta) \}$$

→ Parameters

- $b(y)$  is called the base measure (not depend on  $\eta$ )
- $a(\eta)$  is called the log partition function (not depend on  $y$ )
- $a(\eta)$ ,  $y$  and  $b(y)$  are scalar.  $\eta$  and  $y$  have the same dimensions.



# Example 1: Bernoulli Distribution -> Logistic Regression

Data  $\leftarrow$

$$p(y; \eta) = b(y) \exp \{ \eta^\top y - a(\eta) \}$$

$\rightarrow$  Natural Parameters

## Bernoulli Distribution

$$p(y; \phi) = \phi^y (1 - \phi)^{1-y} = \exp \left\{ y \log \frac{\phi}{1 - \phi} + \log(1 - \phi) \right\}$$

$\searrow \eta$   $\searrow a(\eta)$

Show that term  
is only a function of  $\eta$

# Example 2: Gaussian Distribution -> Linear Regression

Data  $\swarrow$

$$p(y; \eta) = b(y) \exp \{ \eta^\top y - a(\eta) \}$$

$\searrow$  **Natural Parameters**

## Gaussian Distribution

$$p(y; \mu) = \frac{1}{\sqrt{2\pi}} \exp \left\{ -\frac{1}{2}(y - \mu)^2 \right\} = \underbrace{\frac{1}{\sqrt{2\pi}} e^{-y^2/2}}_{b(y)} \left\{ \underbrace{\mu y}_{\eta} - \underbrace{\frac{1}{2}\mu^2}_{a(\eta)} \right\}$$



# Why do we care?

**Data**  $\xleftarrow{\theta^\top x}$

$$p(y; \eta) = b(y) \exp \{ \eta^\top y - a(\eta) \}$$

$\xrightarrow{\text{Natural Parameters}}$

**Inference is Easy:**

$$E[y; \eta] = \frac{da(\eta)}{d\eta}$$

$$\text{Var}[y; \eta] = \frac{d^2 a(\eta)}{d\eta^2}$$

**Learning is Easy:**

**Maximum Likelihood Estimation** leads to **convex** problem in  $\eta$

# Generalized Linear Models

**Assumption:**  $p(y | x; \theta)$  is an exponential family

**Data Type → Probability Distribution**

Binary → Bernoulli → **Logistic Regression**

Real → Gaussian → **Linear Regression**

Counts → Poisson

Positive Real → Gamma, Exponential

Distributions → Dirichlet

# Generalized Linear Models

**Assumption:**  $p(y | x; \theta)$  is an exponential family

The natural parameter is linear in the inputs

$$\eta = \theta^\top x$$

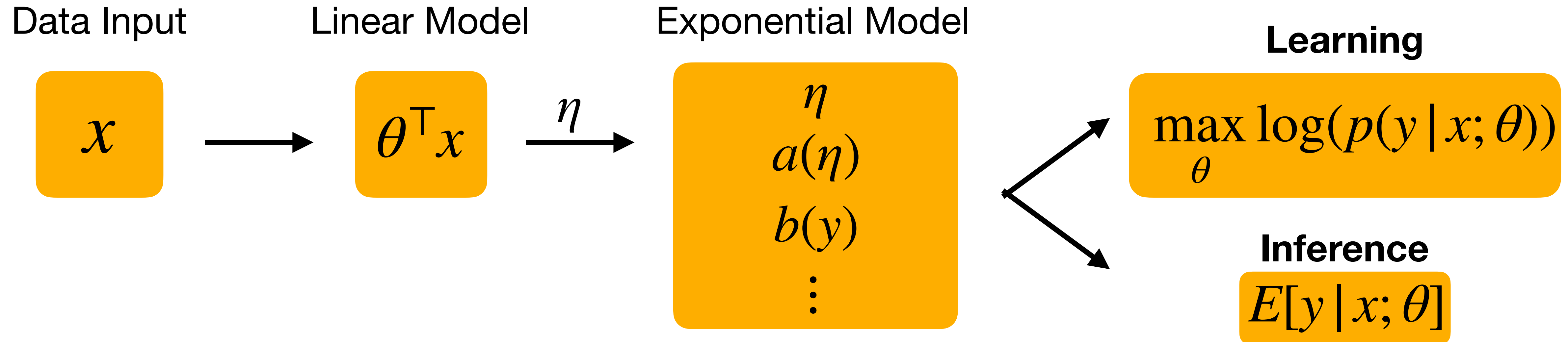
Predictor is a natural consequence

$$h_\theta(x) = E[y | x; \theta]$$



# Generalized Linear Models

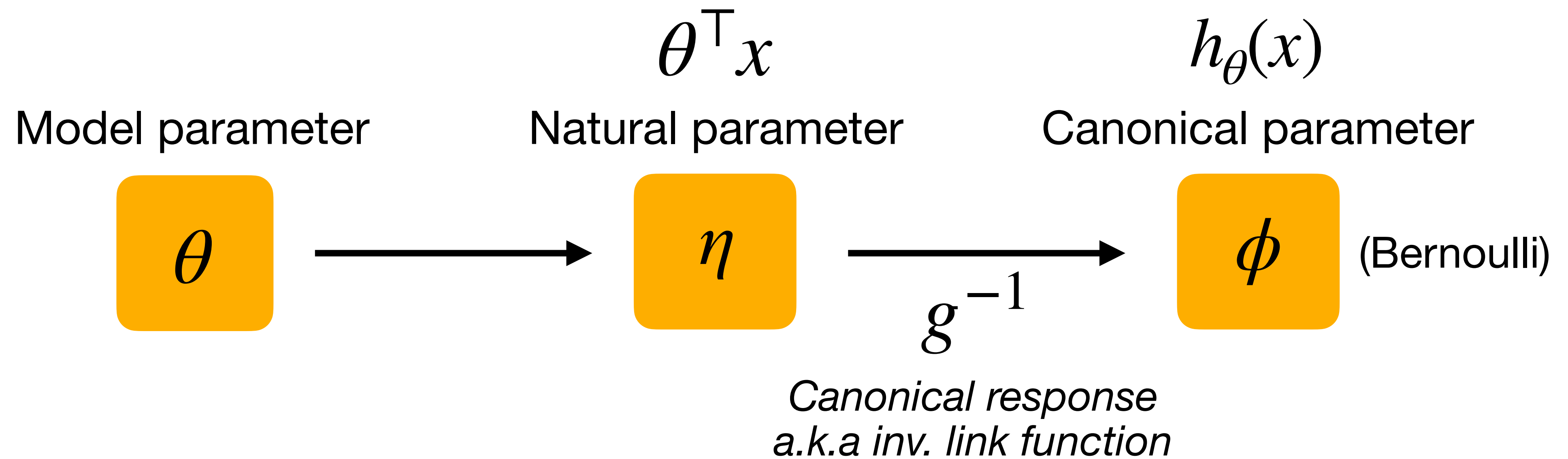
**Assumption:**  $p(y | x; \theta)$  is an exponential family



**Update Rule:**

$$\theta := \theta - \alpha \sum_{i=1}^n \left( h_{\theta} \left( x^{(i)} \right) - y^{(i)} \right) x^{(i)}$$

# Terminology



**Logistic Regression:**

$$h_\theta(x) = E[y | x; \theta]$$

$$\phi = \frac{1}{1 + e^{-\eta}} = \frac{1}{1 + e^{-\theta^\top x}}$$

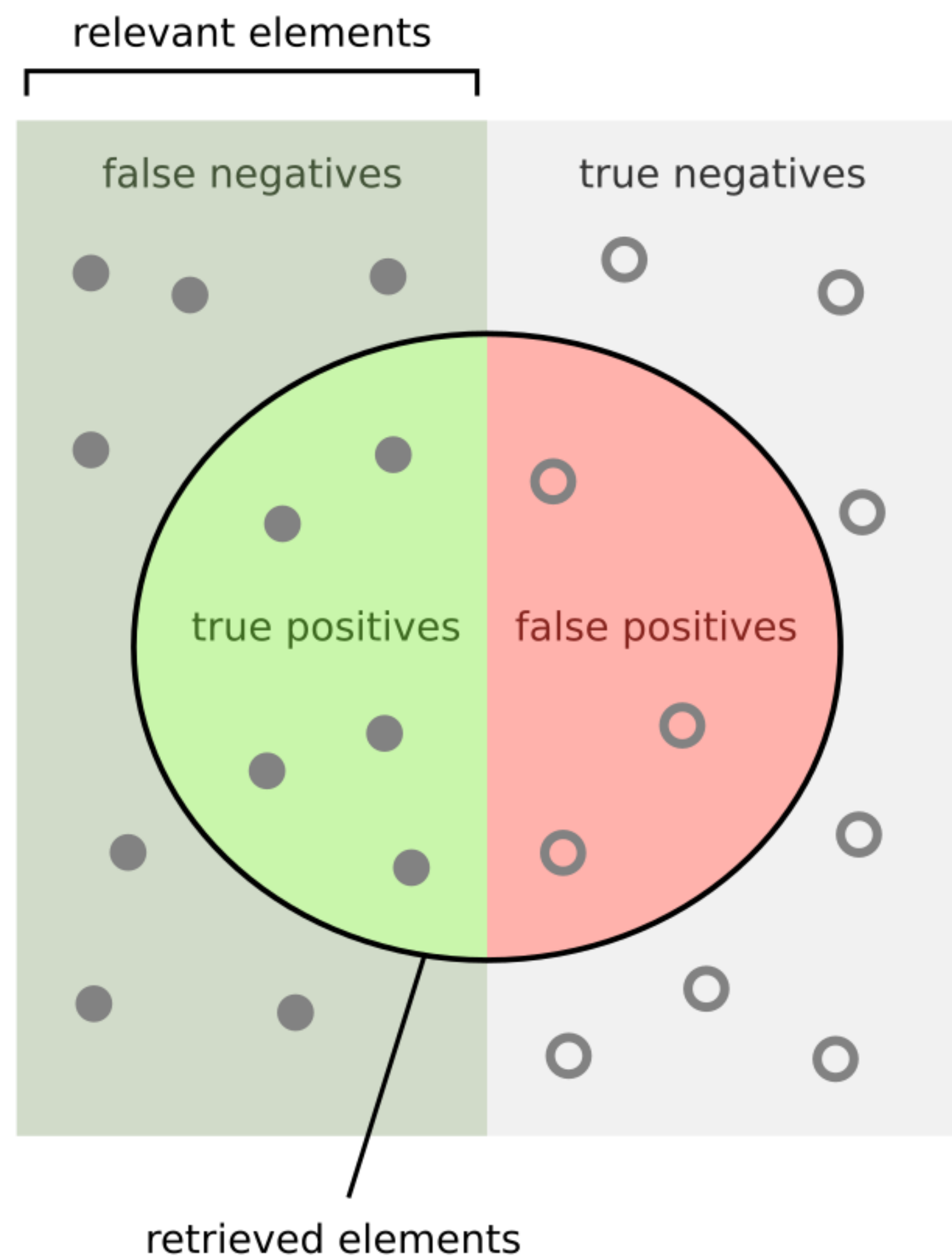
# Evaluation Metrics

Sources: [\[4\]](#)[\[5\]](#)[\[6\]](#)[\[7\]](#)[\[8\]](#)[\[9\]](#)[\[10\]](#)[\[11\]](#) [view](#) • [talk](#) • [edit](#)

|                  |                                | Predicted condition                                                        |                                                                                                                                                           | Sources: <a href="#">[4]</a> <a href="#">[5]</a> <a href="#">[6]</a> <a href="#">[7]</a> <a href="#">[8]</a> <a href="#">[9]</a> <a href="#">[10]</a> <a href="#">[11]</a> <a href="#">view</a> • <a href="#">talk</a> • <a href="#">edit</a> |                                                                                                                                                                                                                        |                                                                                                                                 |
|------------------|--------------------------------|----------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------|
|                  |                                | Total population<br>= P + N                                                | Predicted positive                                                                                                                                        | Predicted negative                                                                                                                                                                                                                            | Informedness, bookmaker<br>informedness (BM)<br>= TPR + TNR − 1                                                                                                                                                        | Prevalence threshold<br>(PT)<br>$= \frac{\sqrt{\text{TPR} \times \text{FPR}} - \text{FPR}}{\text{TPR} - \text{FPR}}$            |
| Actual condition | Positive (P)<br><sup>[a]</sup> | True positive (TP),<br>hit <sup>[b]</sup>                                  | False negative<br>(FN),<br>miss,<br>underestimation                                                                                                       | True positive rate (TPR),<br>recall, sensitivity (SEN),<br>probability of detection, hit rate,<br>power<br>$= \frac{\text{TP}}{\text{P}} = 1 - \text{FNR}$                                                                                    | False negative rate<br>(FNR),<br>miss rate<br>type II error <sup>[c]</sup><br>$= \frac{\text{FN}}{\text{P}} = 1 - \text{TPR}$                                                                                          |                                                                                                                                 |
|                  | Negative (N) <sup>[d]</sup>    | False positive (FP),<br>false alarm, overestimation                        | True negative<br>(TN),<br>correct rejection <sup>[e]</sup>                                                                                                | False positive rate (FPR),<br>probability of false alarm, fall-out<br>type I error <sup>[f]</sup><br>$= \frac{\text{FP}}{\text{N}} = 1 - \text{TNR}$                                                                                          | True negative rate<br>(TNR),<br>specificity (SPC), selectivity<br>$= \frac{\text{TN}}{\text{N}} = 1 - \text{FPR}$                                                                                                      |                                                                                                                                 |
|                  |                                | Prevalence<br>$= \frac{\text{P}}{\text{P} + \text{N}}$                     | Positive predictive value (PPV),<br>precision<br>$= \frac{\text{TP}}{\text{TP} + \text{FP}} = 1 - \text{FDR}$                                             | Negative predictive value<br>(NPV)<br>$= \frac{\text{TN}}{\text{TN} + \text{FN}} = 1 - \text{FOR}$                                                                                                                                            | Positive likelihood ratio<br>(LR+)<br>$= \frac{\text{TPR}}{\text{FPR}}$                                                                                                                                                | Negative likelihood ratio<br>(LR−)<br>$= \frac{\text{FNR}}{\text{TNR}}$                                                         |
|                  |                                | Accuracy<br>(ACC)<br>$= \frac{\text{TP} + \text{TN}}{\text{P} + \text{N}}$ | False discovery rate (FDR)<br>$= \frac{\text{FP}}{\text{TP} + \text{FP}} = 1 - \text{PPV}$                                                                | False omission<br>rate (FOR)<br>$= \frac{\text{FN}}{\text{TN} + \text{FN}} = 1 - \text{NPV}$                                                                                                                                                  | Markedness (MK), deltaP<br>(Δp)<br>= PPV + NPV − 1                                                                                                                                                                     | Diagnostic odds ratio<br>(DOR)<br>$= \frac{\text{LR+}}{\text{LR−}}$                                                             |
|                  |                                | Balanced<br>accuracy (BA)<br>$= \frac{\text{TPR} + \text{TNR}}{2}$         | F <sub>1</sub> score<br>$= \frac{2 \text{ PPV} \times \text{TPR}}{\text{PPV} + \text{TPR}}$ $= \frac{2 \text{ TP}}{2 \text{ TP} + \text{FP} + \text{FN}}$ | Fowlkes–Mallows<br>index (FM)<br>$= \sqrt{\text{PPV} \times \text{TPR}}$                                                                                                                                                                      | <i>phi</i> or Matthews correlation<br>coefficient (MCC)<br>$= \frac{\sqrt{\text{TPR} \times \text{TNR} \times \text{PPV} \times \text{NPV}}}{\sqrt{\text{FNR} \times \text{FPR} \times \text{FOR} \times \text{FDR}}}$ | Threat score (TS), critical<br>success index (CSI),<br>Jaccard index<br>$= \frac{\text{TP}}{\text{TP} + \text{FN} + \text{FP}}$ |



# Precision and Recall



How many retrieved items are relevant?

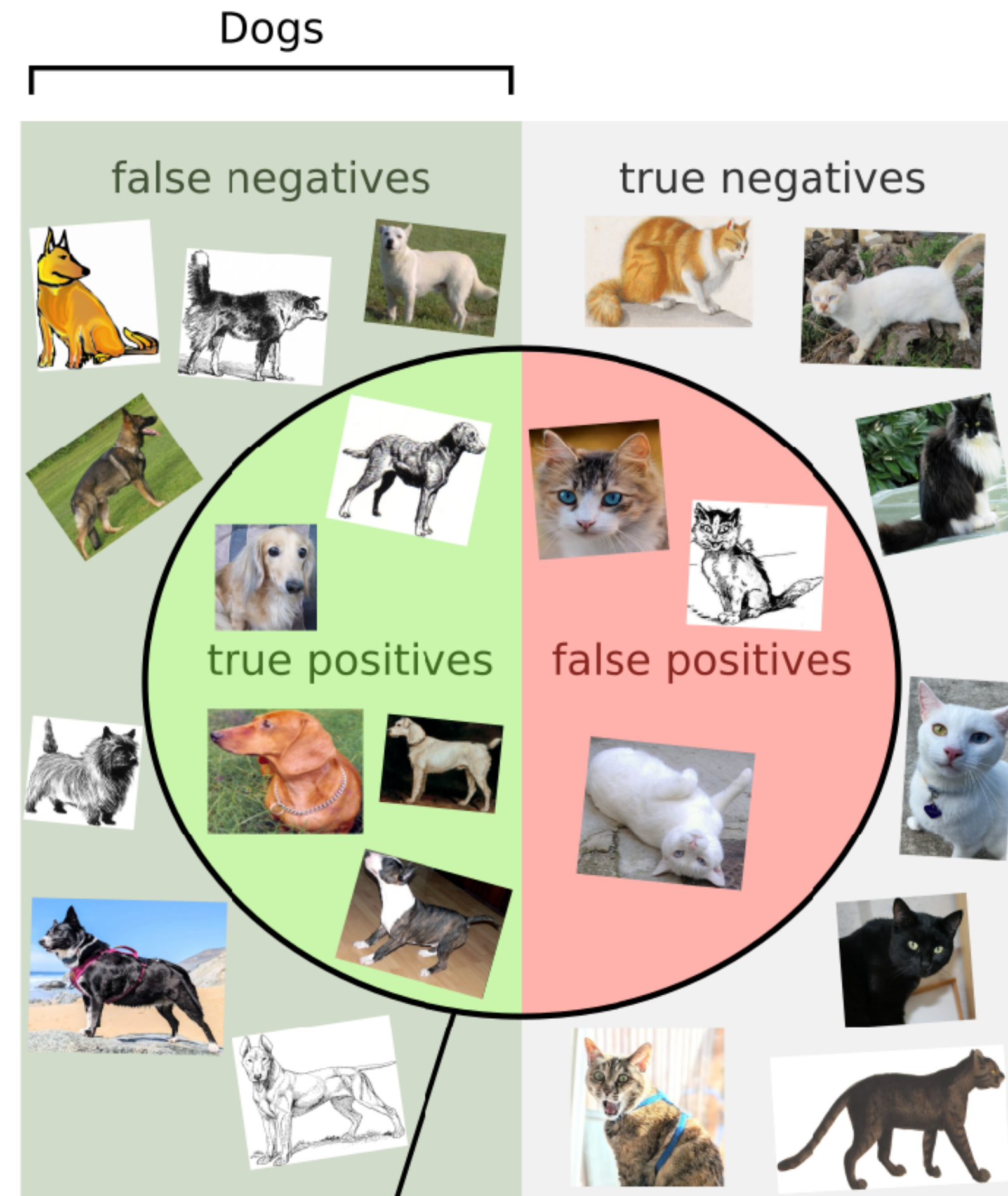
$$\text{Precision} = \frac{\text{true positives}}{\text{true positives} + \text{false positives}}$$

How many relevant items are retrieved?

$$\text{Recall} = \frac{\text{true positives}}{\text{true positives} + \text{false negatives}}$$



# Precision and Recall



$$\text{Precision} = \frac{5 \text{ true pos.}}{8 \text{ total pos.}}$$

$$\text{Recall} = \frac{5 \text{ true pos.}}{12 \text{ total dogs}}$$

$$\text{Prevalence} = \frac{12 \text{ total dogs}}{22 \text{ total images}}$$

$$\text{Accuracy} = \frac{5 \text{ true pos.} + 7 \text{ true neg.}}{22 \text{ total images}}$$

$$F_1 = \frac{2}{\text{recall}^{-1} + \text{precision}^{-1}} = 2 \frac{\text{precision} \cdot \text{recall}}{\text{precision} + \text{recall}} = \frac{2\text{TP}}{2\text{TP} + \text{FP} + \text{FN}}$$

```
sklearn.metrics.f1_score(y_true, y_pred, *, labels=None, pos_label=1,  
average='binary', sample_weight=None, zero_division='warn') # \[source\]
```

Compute the F1 score, also known as balanced F-score or F-measure.

The F1 score can be interpreted as a harmonic mean of the precision and recall, where an F1 score reaches its best value at 1 and worst score at 0. The relative contribution of precision and recall to the F1 score are equal. The formula for the F1 score is:



# Accuracy and Precision

