

Lecture 2 ML & Opt

- Recap of some probability
- Parametrized machine learning models
- Maximum likelihood estimation

Last lecture we saw that we model our ML problem as learning $\rightarrow$ targets

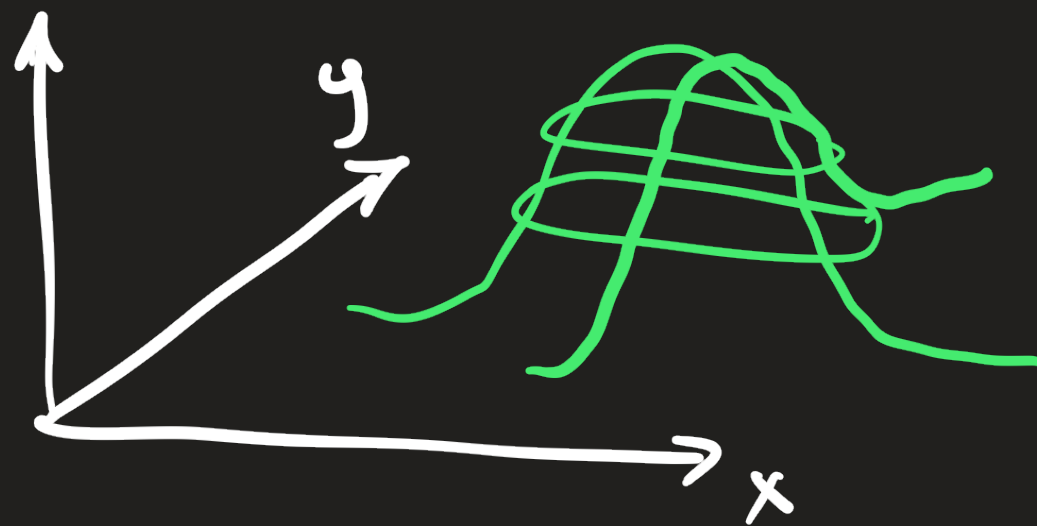
learning

$f: x \mapsto y$

↑
predictor vars / features

→ targets

by modeling $(x, y) \sim P$ and minimizing the average loss of f w.r.t. P



Distributions

$X \sim P$ means $P(X \in A)$ is the probability that X takes on values in the set A

Two cases

X is discrete:

$$P(X \in A) = \sum_{x \in A} p_X(x) \quad \text{where } p_X \text{ is the probability mass function (pmf)}$$

We have $\sum_x p_X(x) = 1$ and $p_X \geq 0$ everywhere

X is continuous:

$$P(X \in A) = \int_A p_X(x) dx \quad \text{where } p_X \text{ is the probability density function (pdf)}$$

We have $\int_{\mathbb{R}} p_X(x) dx = 1$ and $p_X \geq 0$ everywhere

Examples

$X \sim \text{Bern}(p)$ — Bernoulli distribution w/ parameter p
 $p \in [0, 1]$

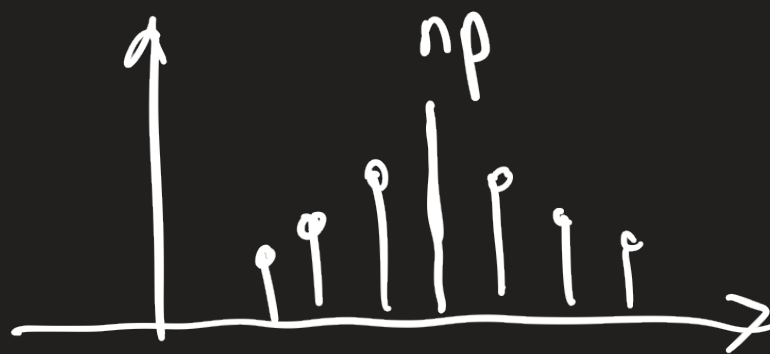
means

$$X = \begin{cases} 1 & \text{with prob } p \\ 0 & \text{w.p. } 1-p \end{cases}$$

$X \sim \text{Binom}(n, p)$ — Binomial dist. w/ params n & p

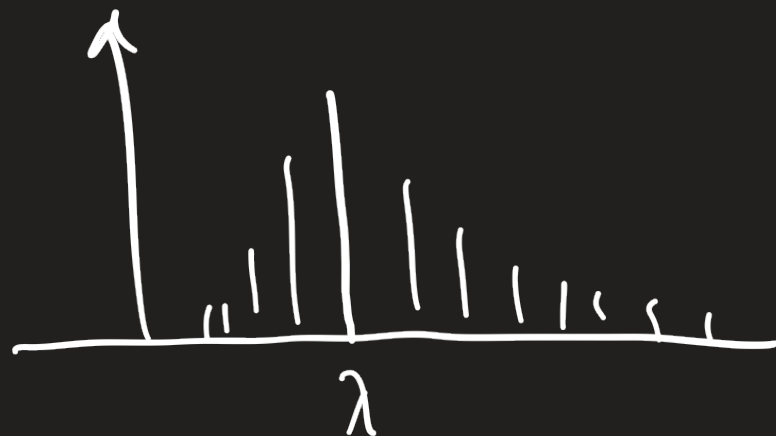
means

$$P(X=k) = \binom{n}{k} p^k (1-p)^{n-k} \text{ for } k \in \{0, 1, \dots, n\}$$



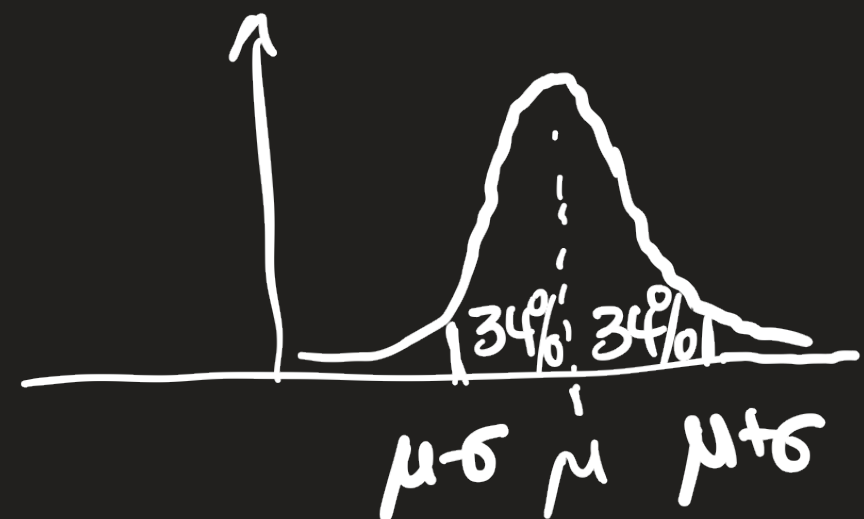
$X \sim \text{Poisson}(\lambda)$ - Poisson dist w/ rate parameter λ

$$P(X=k) = e^{-\lambda} \frac{\lambda^k}{k!} \quad \text{for } k=0,1,2,\dots$$



$X \sim \mathcal{N}(\mu, \sigma^2)$ - Normal or Gaussian r.v. w/
mean μ and variance σ^2

$$p_X(x) = \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{1}{2} \left(\frac{x-\mu}{\sigma} \right)^2}$$



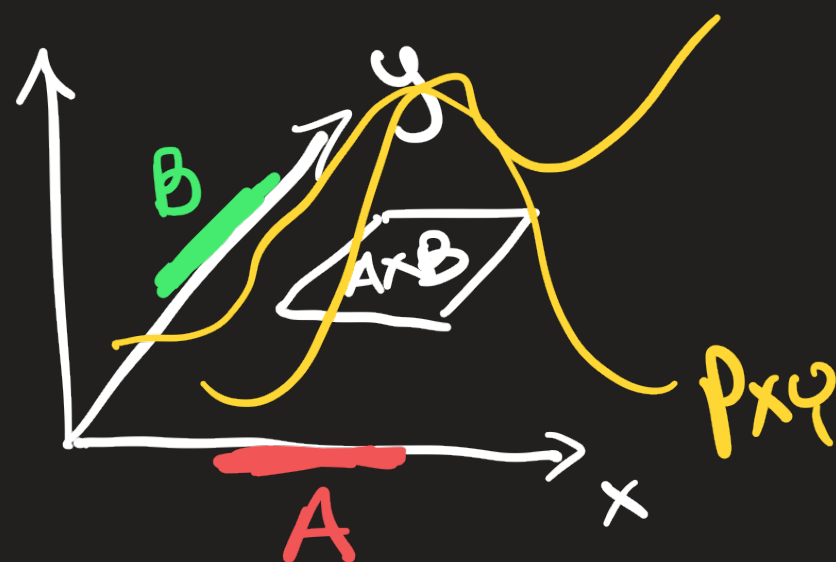
To model dependencies b/w random variables we use joint distributions

$$P(X \in A, Y \in B) \text{ written as } (X, Y) \sim P$$

Analogously we have joint pdfs and joint pmfs.

We assume there exists a joint pdf

$$P(X \in A, Y \in B) = \int_{A \times B} p_{XY}(x, y) dx dy$$



where $p_{XY} \geq 0$ and

$$\int_{\mathbb{R} \times \mathbb{R}} p_{XY} dx dy = 1$$

Marginals

Given $p_{X,Y}$ we can extract the probability dist of X

$$P(X \in A) = P(X \in A, Y \in \mathbb{R}) = \int_{A \times \mathbb{R}} p_{X,Y}(x,y) dx dy$$

$$= \int_A \int_{\mathbb{R}} p_{X,Y}(x,y) dy dx$$

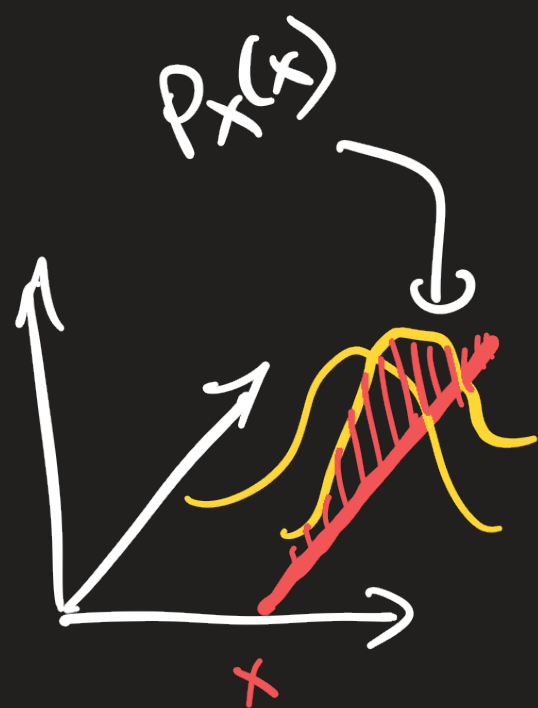
marginal pdf
of x

$$:= \int_A p_X(x) dx$$

where

$$p_X(x) = \int_{\mathbb{R}} p_{X,Y}(x,y) dy$$

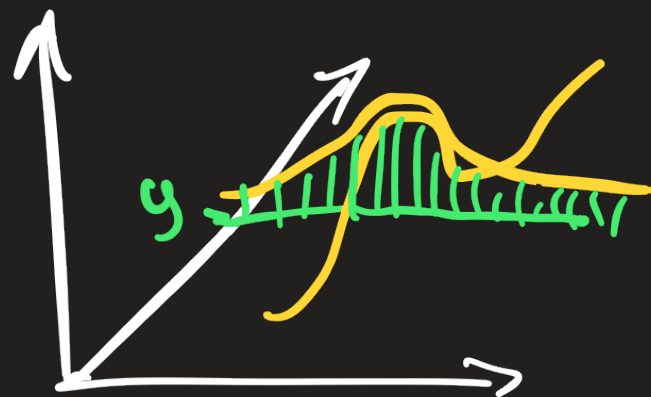
$$\text{and } p_X(x) \geq 0 \text{ and } \int_{\mathbb{R}} p_X(x) dx = 1$$



Similarly we get p_Y by marginalizing over x

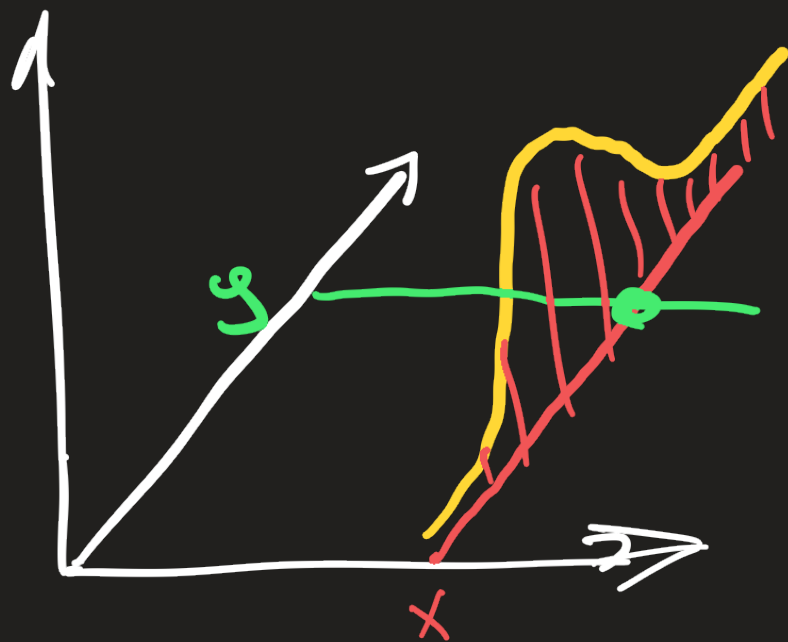
$$p_Y(y) = \int_{\mathbb{R}} p_{XY}(x, y) dx$$

marginal pdf of Y



In the best case, knowledge of y conditional on x is captured by the conditional distribution of y given x

$$P_{Y|X=x}(y) = P_Y(y|X=x) = \frac{P_{X,Y}(x,y)}{P_X(x)}$$



Expectation

Average value of a r.v. X

$$\mathbb{E}X = \int_{\mathbb{R}} x p(x) dx$$

Ex

$$X \sim \text{Bern}(p) \Rightarrow \mathbb{E}X = 0 \cdot (1-p) + 1 \cdot p = p$$

$$X \sim \mathcal{N}(\mu, \sigma^2) \Rightarrow \mathbb{E}X = \mu$$

Expectation of random vectors

$$\text{If } \begin{bmatrix} X \\ Y \end{bmatrix} \sim P \text{ then } \mathbb{E} \begin{bmatrix} X \\ Y \end{bmatrix} = \int_{\mathbb{R} \times \mathbb{R}} \begin{bmatrix} x \\ y \end{bmatrix} p_{X,Y}(x,y) dx dy$$

$$= \int \begin{bmatrix} x p_{X,Y}(x,y) \\ y p_{X,Y}(x,y) \end{bmatrix} dx dy$$

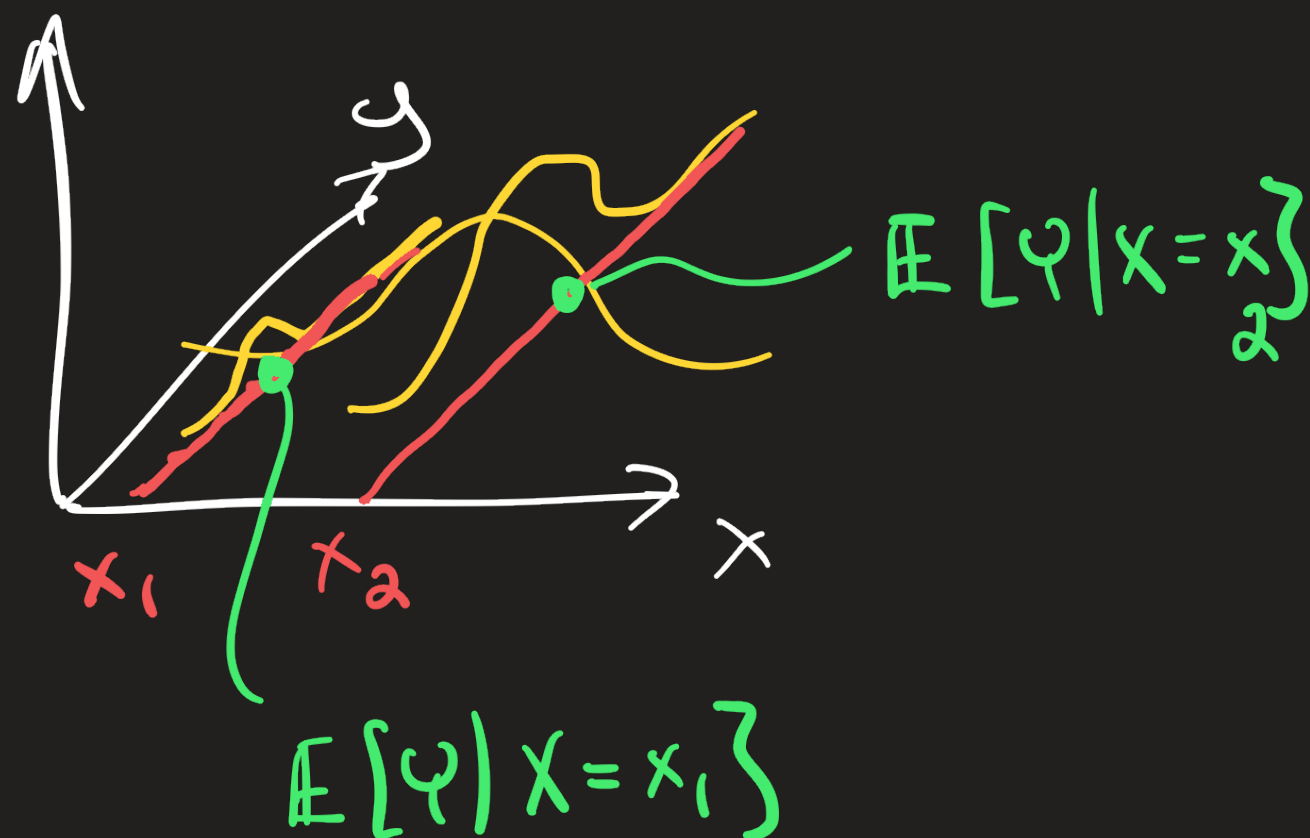
$$= \begin{bmatrix} \int x p_{X,Y}(x,y) dx dy \\ \int y p_{X,Y}(x,y) dx dy \end{bmatrix} \stackrel{\text{verify}}{=} \begin{bmatrix} \int x p_X(x) dx \\ \int y p_Y(y) dy \end{bmatrix}$$

$$= \begin{bmatrix} \mathbb{E} X \\ \mathbb{E} Y \end{bmatrix}$$

Conditional Expectation

$E[Y|X]$ — average value of Y given knowledge of X

$$E[Y|X=x] = \int y p_Y(y|X=x) dy$$



Variance

Variance measures how close a random variable stays to its mean, on average

$$\text{Var}(X) = \mathbb{E}(X - \mathbb{E}X)^2$$

$$= \int_{\mathbb{R}} (x - \mathbb{E}X)^2 p_X(x) dx$$

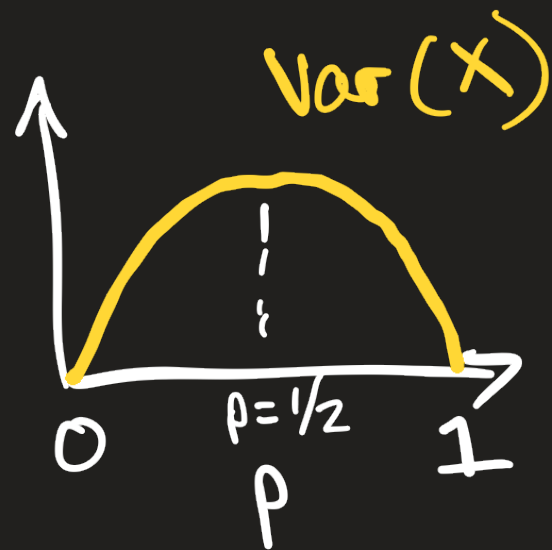
$$= \int_{\mathbb{R}} [x^2 - 2x\mathbb{E}X + (\mathbb{E}X)^2] p_X(x) dx$$

$$= \int_{\mathbb{R}} x^2 p_X(x) dx - 2\mathbb{E}X \int_{\mathbb{R}} x p_X(x) dx + (\mathbb{E}X)^2$$

$$= \mathbb{E}(X^2) - 2(\mathbb{E}X)^2 + (\mathbb{E}X)^2 = \boxed{\mathbb{E}X^2 - (\mathbb{E}X)^2}$$

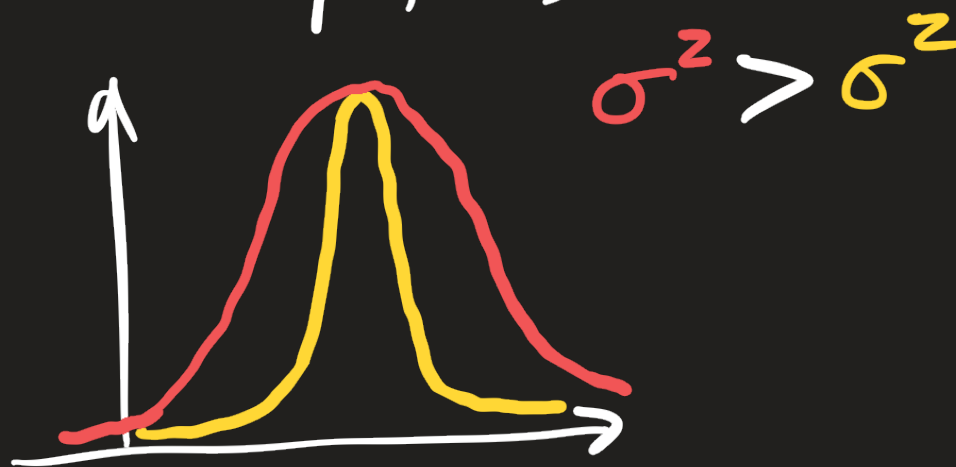
Ex of Variance

$$X \sim \text{Bern}(p) \Rightarrow$$



$$\begin{aligned}\text{Var}(X) &= \mathbb{E}[X^2] - (\mathbb{E}X)^2 \\ &= \mathbb{E}X - (\mathbb{E}X)^2 \\ &= p - p^2\end{aligned}$$

$$X \sim \mathcal{N}(\mu, \sigma^2)$$



Variance is a
measure of
uncertainty

Independence

$$(X \perp\!\!\!\perp Y)$$

We say two random variables are independent if

$$P(X \in A, Y \in B) = P(X \in A) P(Y \in B)$$



$$p_{XY}(x, y) = p_X(x) p_Y(y)$$



$$p_Y(y | X=x) = p_Y(y) \text{ \& } p_X(x | Y=y) = p_X(x)$$

Conditional Independence

We say X & Y are conditionally independent given Z if $(X \perp\!\!\!\perp Y | Z)$

$$P_{X,Y|Z} = P_{X|Z} \cdot P_{Y|Z}$$

$$P_{X,Y|Z=z}(x,y) = \frac{P_{X,Y,Z}(x,y,z)}{P_Z(z)} \quad \text{in general}$$

$$= P_{X|Z=z}(x) P_{Y|Z=z}(y)$$

if $X \perp\!\!\!\perp Y | Z$

Ex:
SAT(A)

GRE(B)

Academic
Ability (Z)

$A \perp\!\!\!\perp B | Z$

Parameterized ML models

$p_Y(y|x)$ — captures everything about values of Y
given knowledge $X=x$

$\mathbb{E}[Y|X=x]$ — captures a point estimate of Y given x

Issues w/ using these in ML: either can be
an arbitrarily complicated unknown function of x

To deal with this, we use parameterized models,
i.e.

$p(y|x) = f_{\theta}(x)$ where $\theta \in \mathbb{R}^P$ is a
parameter vector

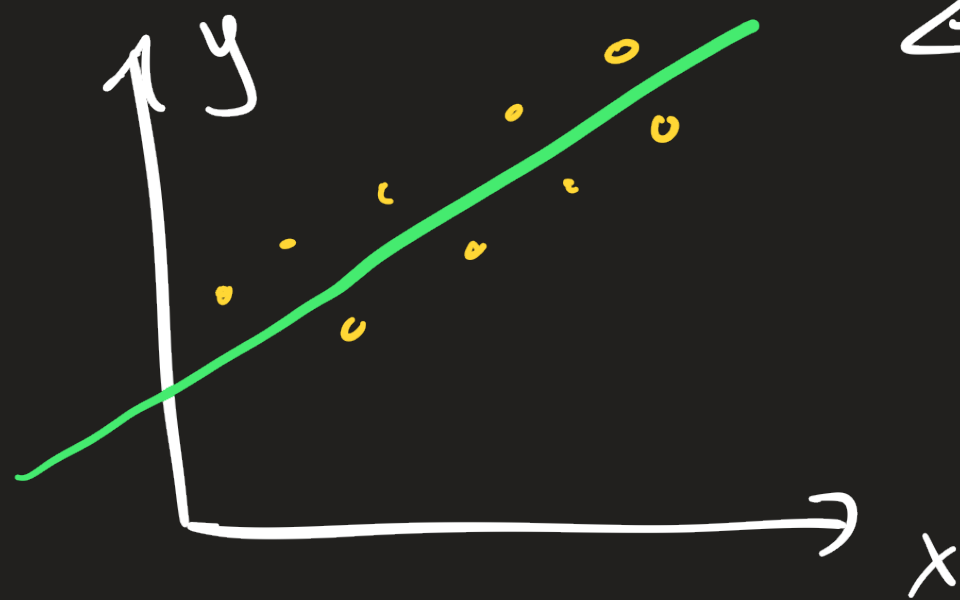
$$= f(x; \theta)$$

Ex

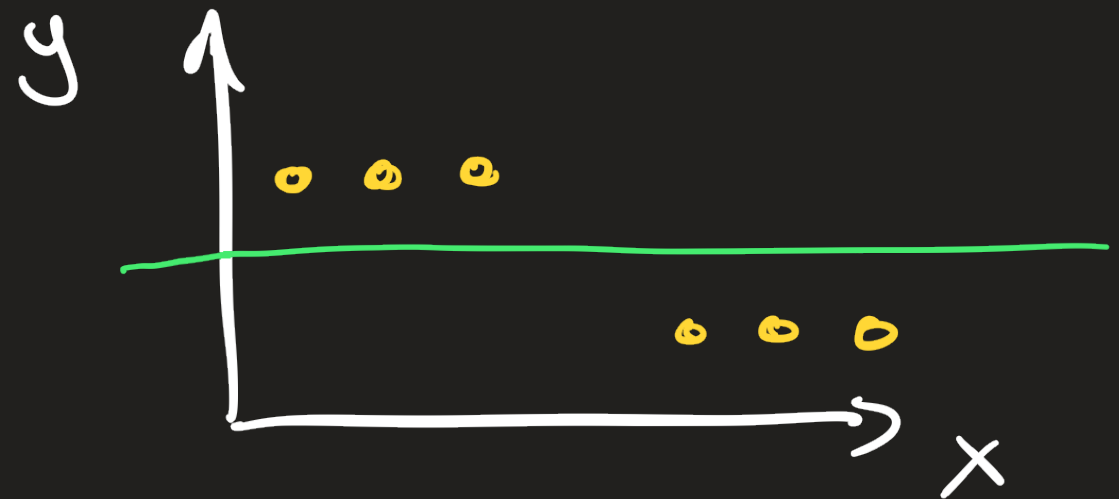
OLS: $y = \beta^T x + \epsilon$ where $\epsilon \sim \mathcal{N}(0, \frac{1}{n})$
i.e. $y|x \sim \mathcal{N}(\beta^T x, \frac{1}{n})$

Assume data (x_i, y_i) come from this type of model for $\forall x$, then we learn β that fits this data best

OLS is appropriate for data that looks like



← and inappropriate for, e.g.



what about parameterized models for data that doesn't fit the OLS assumptions?

e.g.

- y is not a cont. r.v.
- y has a limited range
- $y - E[y|x]$ is not indep of x

Ex

- Poisson / "shot noise" model
corresponds to $\text{Poisson}(\lambda)$ being used, where λ is a function of x

$$y|x \sim \text{Poisson}(\lambda_{\theta}(x))$$

e.g. this gives

$$\mathbb{E}[y|x] = \lambda_{\theta}(x)$$

Notice since $\lambda_{\theta}(x)$ is a rate, it must satisfy $\lambda_{\theta}(x) > 0$. We impose this by taking

$$\lambda_{\theta}(x) = e^{f(x; \theta)}$$