

Probability, Maximum likelihood, and MAP Estimators

Tom Mitchell
Machine Learning 10-601
Jan 26 2009

part of this material is based on slides from Carlos Guestrin, and
from Eric Xing

You should know

- Events
 - discrete random variables, continuous random variables, compound events
- Axioms of probability
 - What defines a reasonable theory of uncertainty
- Independent events
- Conditional probabilities
- Bayes rule and beliefs
- Joint probability distribution

Expected values

Given discrete random variable X , the expected value of X , written $E[X]$ is

$$E[X] = \sum_{x \in \mathcal{X}} x P(X = x)$$

We also can talk about the expected value of functions of X

$$E[f(X)] = \sum_{x \in \mathcal{X}} f(x) P(X = x)$$

Covariance

Given two discrete r.v.'s X and Y , we define the covariance of X and Y as

$$Cov(X, Y) = E[\underbrace{(X - E(X))}_{\text{deviation of } X} \underbrace{(Y - E(Y))}_{\text{deviation of } Y}]]$$

e.g., X =gender, Y =playsFootball ←

or X =gender, Y =leftHanded

Remember:

$$E[X] = \sum_{x \in \mathcal{X}} x P(X = x)$$

Your first consulting job

$$f: x \rightarrow y$$
$$P(Y|X)$$

- A billionaire from the suburbs of Seattle asks you a question:
 - He says: I have thumbtack, if I flip it, what's the probability it will fall with the nail up?
 - You say: Please flip it a few times:



- You say: The probability is: .6
- **He says: Why???**
- You say: Because...

Thumbtack – Binomial Distribution

- $P(\text{Heads}) = \theta$, $P(\text{Tails}) = 1 - \theta$

D : $\downarrow \downarrow \nwarrow \downarrow \nwarrow$
 $X_1 \quad X_2 \quad X_3 \quad X_4 \quad X_5$
 $\theta \times \theta \times (1 - \theta) \times \theta \times (1 - \theta) =$

identically distributed

$$\theta^3 (1 - \theta)^2$$

- Flips are i.i.d.:

- ☐ Independent events
- ☐ Identically distributed according to Binomial distribution

- Sequence D of α_H Heads and α_T Tails

data likelihood $\rightarrow$ $P(D | \theta) = \theta^{\alpha_H} (1 - \theta)^{\alpha_T}$

Maximum Likelihood Estimation

- **Data:** Observed set D of α_H Heads and α_T Tails
- **Hypothesis:** Binomial distribution
- Learning θ is an optimization problem
 - What's the objective function?
- **MLE:** Choose θ that maximizes the probability of observed data:

$$\begin{aligned}\hat{\theta} &= \arg \max_{\theta} P(\mathcal{D} \mid \theta) \\ &= \arg \max_{\theta} \ln P(\mathcal{D} \mid \theta)\end{aligned}$$

Maximum Likelihood Estimate for θ


$$\begin{aligned}\hat{\theta} &= \arg \max_{\theta} \ln P(\mathcal{D} \mid \theta) \\ &= \arg \max_{\theta} \ln \theta^{\alpha_H} (1 - \theta)^{\alpha_T}\end{aligned}$$


- Set derivative to zero:

$$\frac{d}{d\theta} \ln P(\mathcal{D} \mid \theta) = 0$$

■ Set derivative to zero:

$$\frac{d}{d\theta} \ln P(\mathcal{D} | \theta) = 0$$

$$\hat{\theta} = \arg \max_{\theta} \ln P(\mathcal{D} | \theta)$$

$$= \arg \max_{\theta} \ln \theta^{\alpha_H} (1 - \theta)^{\alpha_T} \stackrel{\partial = d}{=} \frac{d}{d\theta} \alpha_H \ln \theta + \alpha_T \ln(1 - \theta)$$

$$\frac{\partial \ln \theta}{\partial \theta} = \frac{1}{\theta}$$

$$\alpha_H \frac{\partial \ln \theta}{\partial \theta} + \alpha_T \frac{\partial \ln(1 - \theta)}{\partial \theta}$$

$$\alpha_H \frac{1}{\theta} + \alpha_T \frac{\frac{\partial(1 - \theta)}{\partial \theta} \cdot \frac{\partial \ln(1 - \theta)}{\partial(1 - \theta)}}{\partial \theta}$$

$$0 = \alpha_H \frac{1}{\theta} + \alpha_T (-1) \frac{1}{1 - \theta}$$

$$\theta = \frac{\alpha_H}{\alpha_H + \alpha_T}$$

How many flips do I need?


$$\hat{\theta}_{MLE} = \frac{\alpha_H}{\alpha_H + \alpha_T}$$

Maximum
likelihood
Estimate.

Bayesian Learning

$$MLE = \underset{\theta}{\operatorname{argmax}} \underbrace{P(\mathcal{D} | \theta)}$$

- Use Bayes rule:

$$\underset{\theta}{\operatorname{argmax}} \underbrace{P(\theta | \mathcal{D})} = \frac{P(\mathcal{D} | \theta) P(\theta)}{P(\mathcal{D})}$$

- Or equivalently:

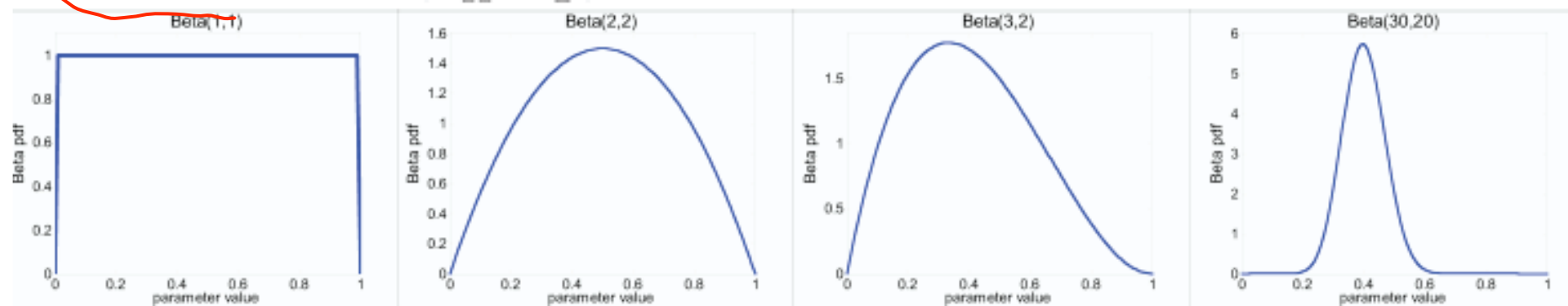
$$\underbrace{P(\theta | \mathcal{D})} \propto P(\mathcal{D} | \theta) P(\theta)$$

Beta prior distribution – $P(\theta)$

$$P(\theta) = \frac{\theta^{\beta_H-1}(1-\theta)^{\beta_T-1}}{B(\beta_H, \beta_T)} \sim \text{Beta}(\beta_H, \beta_T)$$

Mean:

Mode:



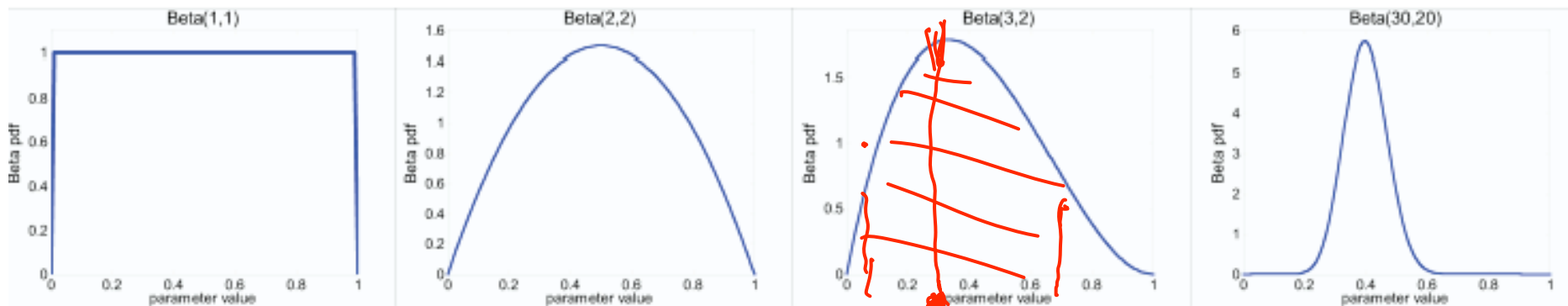
- Likelihood function: $P(\mathcal{D} | \theta) = \theta^{\alpha_H}(1 - \theta)^{\alpha_T}$
- Posterior: $P(\theta | \mathcal{D}) \propto P(\mathcal{D} | \theta)P(\theta)$

Posterior distribution

- Prior: $Beta(\beta_H, \beta_T)$
- Data: α_H heads and α_T tails

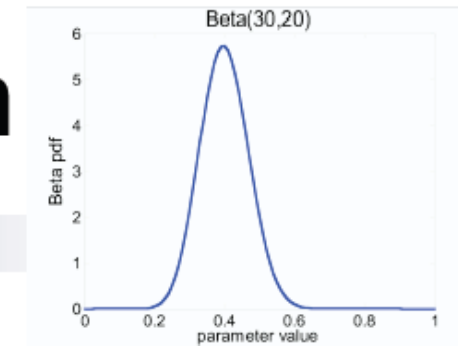
- Posterior distribution:

$$P(\theta \mid \mathcal{D}) \sim Beta(\beta_H + \alpha_H, \beta_T + \alpha_T)$$



MAD est

MAP for Beta distribution



$$P(\theta \mid \mathcal{D}) = \frac{\theta^{\beta_H + \alpha_H - 1} (1 - \theta)^{\beta_T + \alpha_T - 1}}{B(\beta_H + \alpha_H, \beta_T + \alpha_T)} \sim \text{Beta}(\beta_H + \alpha_H, \beta_T + \alpha_T)$$

- MAP: use most likely parameter:

$$\hat{\theta} = \arg \max_{\theta} P(\theta \mid \mathcal{D}) = \frac{\alpha_H + \beta_H}{\alpha_H + \alpha_T + \beta_H + \beta_T}$$

- Beta prior equivalent to extra thumbtack flips
- As $N \rightarrow \infty$, prior is “forgotten”
- **But, for small sample size, prior is important!**

Dirichlet distribution

- number of heads in N flips of a two-sided coin
 - follows a binomial distribution
 - Beta is a good prior (conjugate prior for binomial)
- what it's not two-sided, but k-sided?
 - follows a multinomial distribution
 - Dirichlet distribution is the conjugate prior

$$P(\theta_1, \theta_2, \dots, \theta_K) = \frac{1}{B(\alpha)} \prod_i^K \theta_i^{(\alpha_i - 1)}$$

Lejeune Dirichlet



Johann Peter Gustav Lejeune Dirichlet

Born	13 February 1805 Düren, French Empire
Died	5 May 1859 (aged 54) Göttingen, Hanover
Residence	 Germany
Nationality	 German
Fields	Mathematician
Institutions	University of Berlin University of Breslau University of Göttingen
Alma mater	University of Bonn
Doctoral advisor	Simeon Poisson Joseph Fourier
Doctoral students	Ferdinand Eisenstein Leopold Kronecker Rudolf Lipschitz Carl Wilhelm Borchardt
Known for	Dirichlet function Dirichlet eta function

Estimating Parameters

- Maximum Likelihood Estimate (MLE): choose θ that maximizes probability of observed data $\mathcal{D}$

$$\hat{\theta} = \arg \max_{\theta} P(\mathcal{D} \mid \theta)$$

- Maximum a Posteriori (MAP) estimate: choose θ that is most probable given prior probability and the data

$$\begin{aligned}\hat{\theta} &= \arg \max_{\theta} P(\theta \mid \mathcal{D}) \\ &= \arg \max_{\theta} = \frac{P(\mathcal{D} \mid \theta)P(\theta)}{P(\mathcal{D})}\end{aligned}$$

You should know

- Probability basics
 - random variables, events, sample space, conditional probs, ...
 - independence of random variables
 - Bayes rule
 - Joint probability distributions
 - calculating probabilities from the joint distribution
- Point estimation
 - maximum likelihood estimates
 - maximum a posteriori estimates
 - distributions – binomial, Beta, Dirichlet, ...



Example: Bernoulli model

- Data:

- We observed N iid coin tossing: $D=\{1, 0, 1, \dots, 0\}$

- Representation:

Binary r.v:

$$x_n = \{0,1\}$$

- Model:

$$P(x) = \begin{cases} 1-\theta & \text{for } x=0 \\ \theta & \text{for } x=1 \end{cases} \Rightarrow P(x) = \theta^x (1-\theta)^{1-x}$$

- How to write the likelihood of a single observation x_i ?

$$P(x_i) = \theta^{x_i} (1-\theta)^{1-x_i}$$

- The likelihood of dataset $D=\{x_1, \dots, x_N\}$:

$$P(x_1, x_2, \dots, x_N | \theta) = \prod_{i=1}^N P(x_i | \theta) = \prod_{i=1}^N (\theta^{x_i} (1-\theta)^{1-x_i}) = \theta^{\sum_{i=1}^N x_i} (1-\theta)^{\sum_{i=1}^N 1-x_i} = \theta^{\text{\#head}} (1-\theta)^{\text{\#tails}}$$



Conditional Independence

Definition: X is conditionally independent of Y given Z, if the probability distribution governing X is independent of the value of Y, given the value of Z

$$(\forall i, j, k) P(X = x_i | Y = y_j, Z = z_k) = P(X = x_i | Z = z_k)$$

Which we often write

$$P(X | Y, Z) = P(X | Z)$$

E.g.,

$$P(\textit{Thunder} | \textit{Rain}, \textit{Lightning}) = P(\textit{Thunder} | \textit{Lightning})$$