

CSC421 Intro to Artificial Intelligence



UNIT 31: Statistical Learning

Outline



- Bayesian Learning
- Maximum a posteriori (MAP) and Maximum Likelihood Learning
- ML parameter learning with complete data

Full Bayesian Learning



- View learning as bayesian updating of a probability distribution over the **hypothesis space**
- **H** is the hypothesis variable, $h_1, \dots, h_n$, prior **$P(H)$**
- **jth** observation gives outcome of rv **D_j**
 - Training data **$\underline{d} = d_1, \dots, d_n$**
- Given the data so far, each hypothesis has a posterior probability
 - $P(h_i | \underline{d}) = ?$

Full Bayesian Learning



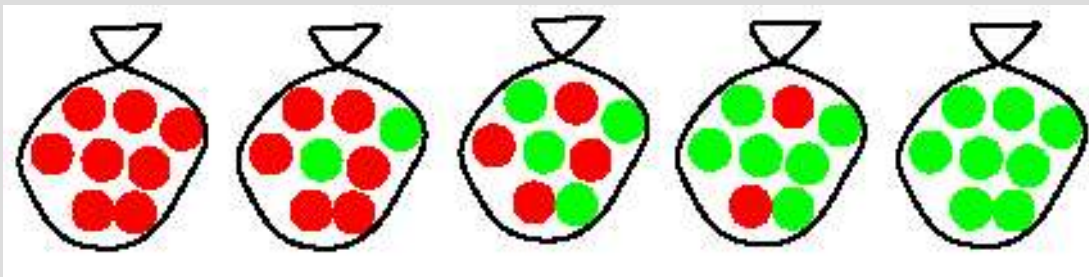
- $P(h_i | \underline{d}) = \frac{1}{Z} P(\underline{d} | h_i) P(h_i)$ where $P(\underline{d} | h_i)$ is called the **likelihood**
- Predictions use a likelihood weighted average over the hypotheses:
 - $P(X | \underline{d}) = \sum_i P(X | \underline{d}, h_i) P(h_i | \underline{d}) = \sum_i P(X | h_i) P(h_i | \underline{d})$
 - Assume each hypothesis determines a probability distribution over X
 - No need to pick best guess
- IID (independently and identically distributed):
 - $P(\underline{d} | h_i) = \prod_j P(d_j | h_i)$

Example



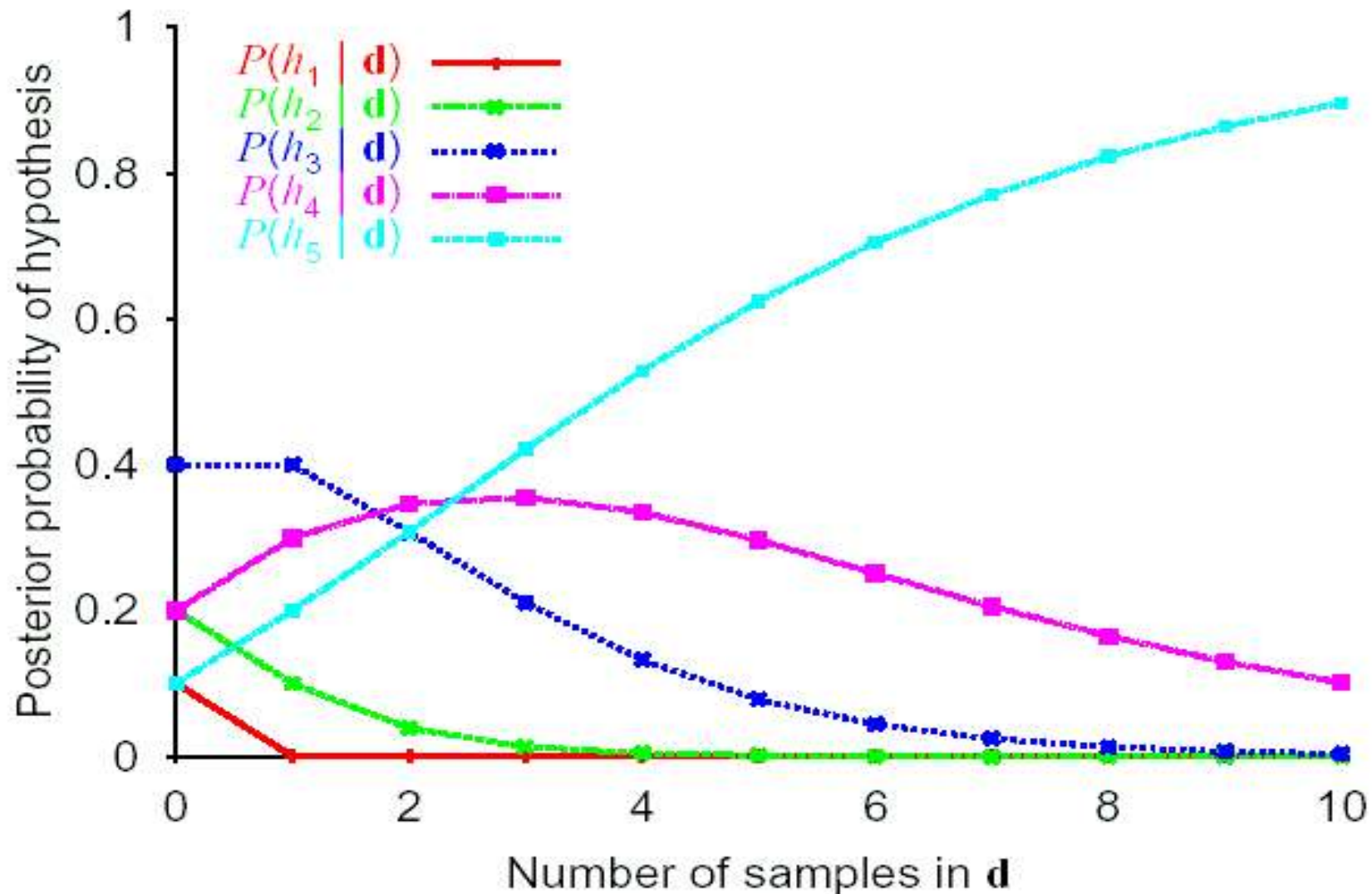
- Suppose there are 5 kinds of bags of candies :
 - 10% are h_1 : 100% cherry candies
 - 20% are h_2 : 75% cherry candies 25% lime
 - 40% are h_3 : 50% cherry 50% lime
 - 20% are h_4 : 25% cherry 75% lime
 - 10% are h_5 : 50% cherry 50% lime

- Observe



What kind of bag is it ?
What flavour will be
the next candy ?

Posterior Probabilities of Hypotheses



MAP approximation

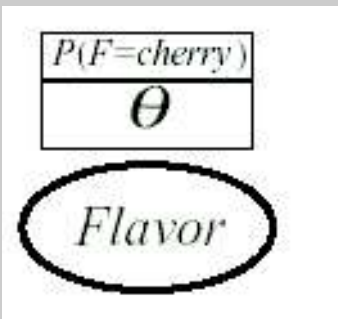


- Summing over entire hypothesis space intractable
- Maximum a Posteriori (MAP) learning: choose h_{MAP} maximizing $P(h_i | \underline{d})$
 - i.e. Maximize $P(\underline{d}|h_i) P(h_i)$ or $\log P(\underline{d}|h_i) + \log P(h_i)$
 - Log terms can be viewed as (negative of)
 - Bits to encode data/hypothesis + bits to encode hypothesis

ML approximation



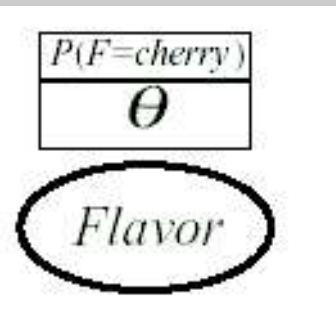
- For large datasets prior becomes irrelevant
- Maximum likelihood (ML) learning: choose h_{ML} maximizing $P(\underline{d}|h_i)$
 - i.e. Simply get the best fit for the data – identical to MAP with uniform prior
- ML is the “standard” (non-Baysian) statistical learning method



ML parameter learning in Bayes nets



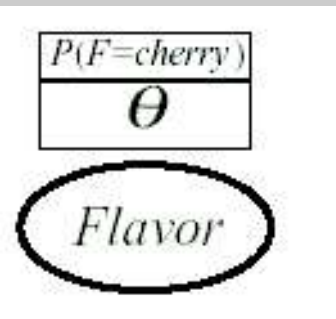
- Bag from a new manufacturer; fraction θ of cherry candies ? Any θ is possible – continuum of hypotheses – θ is a parameter for this simple (binomial) family of models
- Suppose we unwrap N candies, c cherries and $l = N - c$ limes. These are IID:
 - $P(\underline{d} | h_{\theta}) = ?$



ML parameter learning in Bayes nets



- Suppose we unwrap N candies, c cherries and $l = N - c$ limes. These are IID:
 - $P(\underline{d} | h_\theta) = \prod P(d_j | h_\theta) = \theta^c (1-\theta)^l$
- Maximize with respect to θ log likelihood
 - $L(\underline{d} | \theta) = \log P(\underline{d} | h_\theta) = \sum \log P(d_j | h_\theta) = c \log \theta + l \log(1-\theta)$
 - $dL(\underline{d} | \theta) / d\theta = ?$

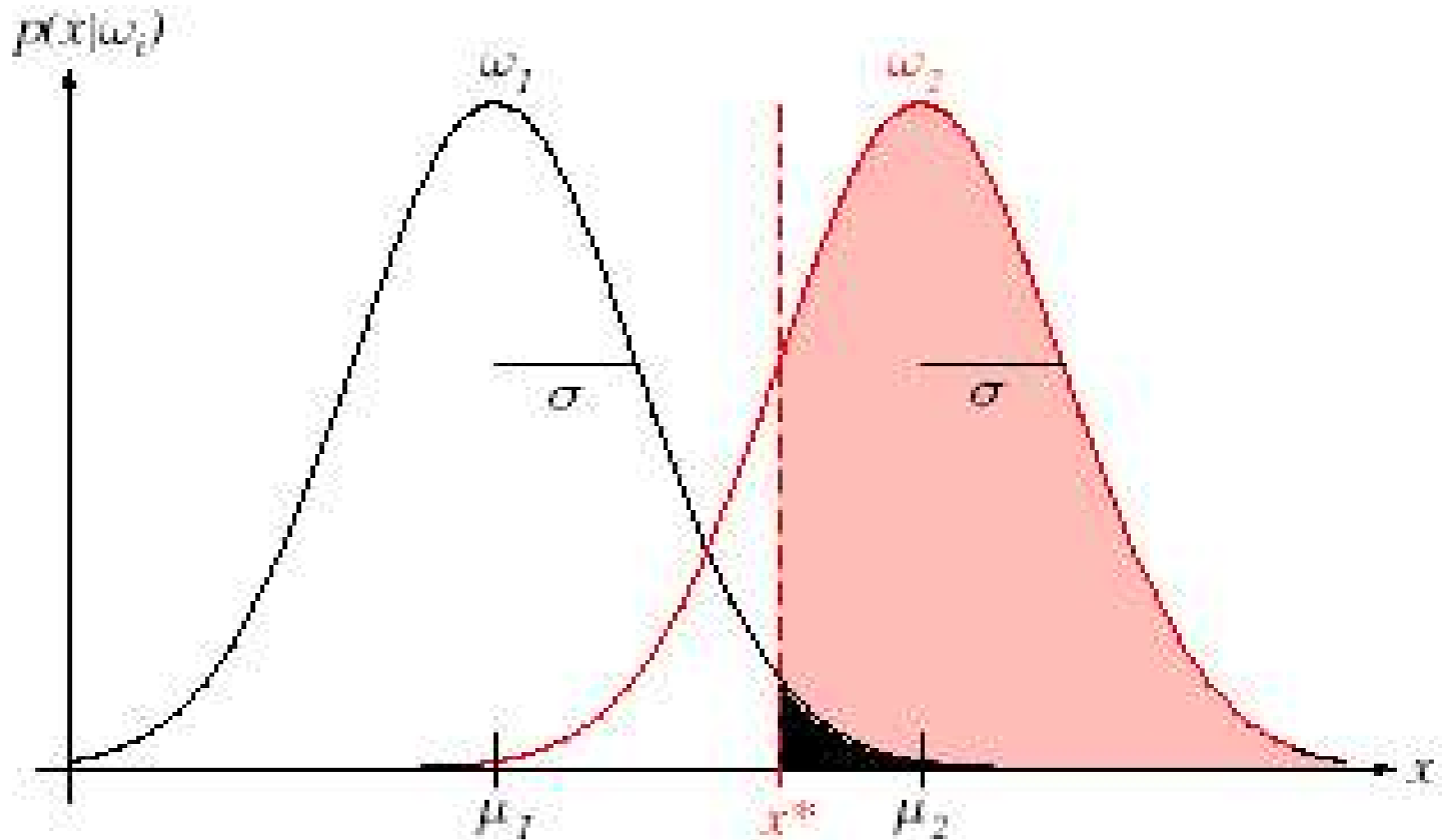


ML parameter learning in Bayes nets



- Suppose we unwrap N candies, c cherries and $l = N - c$ limes. These are IID:
 - $P(\underline{d} | h_{\theta}) = \prod P(d_j | h_{\theta}) = \theta^c (1-\theta)^l$
- Maximize with respect to θ log likelihood
 - $L(\underline{d} | \theta) = \log P(\underline{d} | h_{\theta}) = \sum \log P(d_j | h_{\theta}) = c \log \theta + l \log(1-\theta)$
 - $dL(\underline{d} | \theta) / d\theta = c/\theta - l/(1-\theta) = 0$
 - $\Rightarrow \theta = c / (c + l) = c / N$

Classification using single Gaussian



ML Summary



- ML steps
 - Choose parametrized family of functions
 - Write down likelihood of data as a function of parameters
 - Write derivative of log likelihood
 - Find parameters θ such that derivative is zero
 - Analytically or numerically (optimization)