

Bayesian Statistics

In the previous lectures, especially on maximum Likelihood method, we focused on $P(D|O)$ but noted

that, sometimes, when noise effects are large, it may be better not to take the max likelihood estimate. In other words, we can accept some amount of bias to reduce the impact of variance.

Our ^{extreme} example: $X_i \text{ iid } N(\mu, \sigma^2)$
 $i=1, \dots, N$

MLE for μ : $\hat{\mu}_{MLE} = \bar{X}$

$$E(\bar{X} - \mu)^2 = \frac{\sigma^2}{N}$$

$$MSE_{MLE} = \frac{\sigma^2}{N}$$

No bias, all variance.

However, if $|\mu| < \frac{\sigma}{\sqrt{N}}$ (Sample size too small)

$\hat{\mu} = 0$ has less error.

No variance, but bias.

$$MSE_0 = \mu^2$$

$$MSE_0 < MSE_{MLE}$$

Bayesian approach is one way of balancing ~~bias~~ bias-variance trade-off.

Now we allow talking about probability distribution of θ 's themselves

$$\frac{P(D|\theta) p_0(\theta)}{\int d\theta' P(D|\theta') p_0(\theta')} = P(\theta|D)$$

$\nwarrow$ Prior $\nwarrow$ Posterior
 $\nwarrow$ $P(D)$

Note that $\log P(\theta|D)$

$$= \log P(D|\theta) + \log p_0(\theta) - \log P(D)$$

$\uparrow$ $\uparrow$ $\uparrow$
 log likelihood additional regularizer / penalty function for θ constant

If we want to find the ~~the~~ "most likely" θ , we optimize $\log P(\theta|D)$. The additional regularizer / penalty alters $\hat{\theta}$ from $\hat{\theta}_{MLE}$.

This estimate is called the maximum a posteriori (MAP) estimate.

Several non-Bayesian approaches also advocate using additional penalties for estimation in order to control.

Bayesians ~~believe~~ ^(believe) there is more to Bayesian approach. Before we go there, let us look at some examples, though.

I will skip examples of discrete choice of hypotheses, as discussed in the beginning of Chapter 3. Let us start with the beta-binomial model.

$$X_i \sim \text{Ber}(\theta)$$

$$\text{Prob}[X_i=1] = \theta$$

$$\text{Prob}[X_i=0] = 1-\theta$$

X_i 's independent.

[Ex: Head-tails for a possibly biased coin]

$$P(\mathcal{D}|\theta) = \theta^{N_1} (1-\theta)^{N_0}$$

$$\blacksquare N_1 = \sum_{i=1}^N I(x_i = 1)$$

$$N_0 = \sum_{i=1}^N I(x_i = 0)$$

$$N_1 + N_0 = N$$

Note that N_1 is a sufficient statistics for θ . We don't need to know which order successes happened to infer θ .

$$N_1 \sim \text{Bin}(N, \theta)$$

$$\text{Bin}(k|N, \theta) = \binom{N}{k} \theta^k (1-\theta)^{N-k}$$

Now let us choose a prior. Note that if we choose prior $p(\theta) \propto \theta^{a-1} (1-\theta)^{b-1}$, namely $\theta \sim \text{Beta}(a, b)$, then,

$$P(\theta|\mathcal{D}) \propto P(\mathcal{D}|\theta) p_0(\theta)$$

$$\propto \theta^{N_1+a-1} (1-\theta)^{N_0+b-1}$$

So the posterior distribution follows a beta

distribution as well.

When the posterior and prior belongs to the same family of distributions the prior and the posterior are called conjugate distribution and the prior is called a conjugate prior for the likelihood function.

In this case, the ~~posterior~~ prior is effectively adding some pseudo-counts to N_1, N_0 .

The parameters a, b are called hyperparameters. We will later talk about how to choose these hyperparameters.

If we set $a=1, b=1$, we get an uniform prior.

In principle we could update things sequentially.

prior $\xrightarrow{\text{data 1}}$ Posterior = new prior $\xrightarrow{\text{data 2}}$ new posterior
 is equivalent to

Prior $\xrightarrow{(\text{data 1}, \text{data 2})}$ new posterior

In many circumstances, when data arrives sequentially, you can have an ~~easy~~ online algorithm updating things sequentially.

Posterior mean and mode

Mode $\hat{\theta}_{\text{map}} = \arg \max_{\theta} \log \left[\theta^{N_1+a-1} (1-\theta)^{N_0+b-1} \right]$

θ deriv $\frac{N_1+a-1}{\theta} - \frac{N_0+b-1}{1-\theta} = 0$

$\Rightarrow \frac{\theta}{1-\theta} = \frac{N_1+a-1}{N_0+b-1} \Rightarrow \hat{\theta}_{\text{map}} = \frac{N_1+a-1}{N_0+N_1+a+b-2}$
 $= \frac{N_1+a-1}{N+a+b-2}$

When $a=b=1$ $\hat{\theta}_{\text{map}} = \frac{N_1}{N} = \theta_{\text{MLE}}$

What is the average posterior θ ?

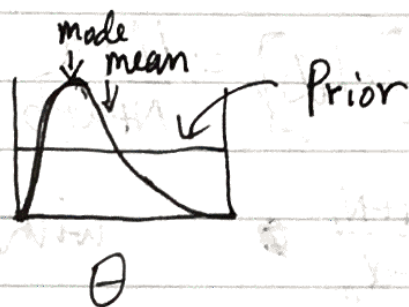
$$\bar{\theta} = \frac{\int_0^1 d\theta \theta^{r_1+1} (1-\theta)^{r_2}}{\int_0^1 d\theta \theta^{r_1} (1-\theta)^{r_2}}$$

where $r_1 = N_1 + a - 1$ $r_2 = N_2 + b - 1$

$$\begin{aligned}\bar{\theta} &= \frac{\int_0^1 d\theta \theta^{r_1+1} (1-\theta)^{r_2}}{\int_0^1 d\theta \theta^{r_1} (1-\theta)^{r_2}} = \frac{B(r_1+2, r_2+1)}{B(r_1+1, r_2+1)} \\ &= \frac{\Gamma(r_1+2) \Gamma(r_2+1) / \Gamma(r_1+r_2+3)}{\Gamma(r_1+1) \Gamma(r_2+1) / \Gamma(r_1+r_2+2)} \\ &= \frac{r_1+1}{r_1+r_2+2} = \frac{N_1+a}{N_1+a+b}\end{aligned}$$

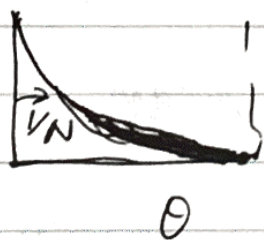
This is different from $\hat{\theta}_{\text{MAP}}$
Even if I took $a=b=1$,
I get $\bar{\theta} = \frac{N_1+1}{N+2}$

The mean posterior seems to have inserted 'pseudo counts' by itself. This is an example of Bayesian averaging 'regularizing' the estimate



Such ~~smooth~~ 'smoothing' is important when data is sparse. Imagine $N_1 = 0$, no success observed so far.

Can I conclude $\theta = 0$? With uniform prior. $P(\theta|D) \propto (1-\theta)^N \approx e^{-N\theta}$



Intuitively that observation is saying that we believe $\theta \sim O(1/N)$

and that is why perhaps we have not observed it.

$$\bar{\theta} = \frac{0+1}{N+2} = \frac{1}{N+2}$$

Posterior mean $\frac{N+1}{N+2}$ as a estimate gives

Laplace's rule of succession
and it is also known as
add-one smoothing!

Note that $E[\theta | \mathcal{D}] = \frac{N_1 + a}{N + a + b}$

$$= \frac{\alpha_0 + N_1}{N + \alpha_0}$$

with $\alpha_0 = a + b$

$$m_1 = \frac{a}{a + b}$$

↑
Prior mean



$$E[\theta | \mathcal{D}] = \underbrace{\left(\frac{\alpha_0}{N + \alpha_0} \right)}_{\text{Prior mean}} m_1 + \underbrace{\left(\frac{N}{N + \alpha_0} \right)}_{\text{MLE}} \frac{N_1}{N}$$

$$= \lambda m_1 + (1 - \lambda) \hat{\theta}_{MLE}$$

↑ Convex combination

As $N \rightarrow \infty$ $\hat{\theta}$ goes to $\hat{\theta}_{MLE}$

With lots of data, priors do not matter.

Posterior Variance

$$\text{var}[\theta | D] = E[\theta^2 | D] - E[\theta | D]^2$$

$$= \frac{\Gamma(\gamma_1+3)\Gamma(\gamma_2+1)/\Gamma(\gamma_1+\gamma_2+4)}{\Gamma(\gamma_1+1)\Gamma(\gamma_2+1)/\Gamma(\gamma_1+\gamma_2+2)} - \left(\frac{\gamma_1+1}{\gamma_1+\gamma_2+2}\right)^2$$

$$= \frac{(\gamma_1+2)(\gamma_1+1)}{(\gamma_1+\gamma_2+3)(\gamma_1+\gamma_2+2)} - \frac{(\gamma_1+1)(\gamma_1+1)}{(\gamma_1+\gamma_2+2)(\gamma_1+\gamma_2+2)}$$

$$= \frac{\gamma_1+1}{\gamma_1+\gamma_2+2} \left[\frac{\gamma_1+2}{\gamma_1+\gamma_2+3} - \frac{\gamma_1+1}{\gamma_1+\gamma_2+2} \right]$$

$$= \frac{(\gamma_1+1) [2(\gamma_1+\gamma_2)+4 - \{4\gamma_1+\gamma_2+3\}]}{(\gamma_1+\gamma_2+2)^2(\gamma_1+\gamma_2+3)}$$

$$= \frac{(\gamma_1+1)(\gamma_2+1)}{(\gamma_1+\gamma_2+2)^2(\gamma_1+\gamma_2+3)}$$

$$= \frac{(N_1+a)(N_0+b)}{(N_1+N_0+a+b)^2(N_1+N_0+a+b+1)}$$

~~$$\frac{(N_1+a)(N_0+b)}{(N_1+N_0+a+b)^2(N_1+N_0+a+b+1)}$$~~

~~$$\frac{(N_1+a)(N_0+b)}{(N_1+N_0+a+b)^2(N_1+N_0+a+b+1)}$$~~

When N is large compared to a, b

$$\boxed{\text{Var}[\theta|D]} \approx \frac{N_1 N_0}{N^3}$$

$$\approx \frac{\hat{\theta}(1-\hat{\theta})}{N}$$

where $\hat{\theta} = \hat{\theta}_{MLE}$.

This is related to $\hat{\theta}_{MLE} = \frac{N_1}{N}$
following central limit theorem
with the std / error bar

$$\sigma = \sqrt{\frac{\hat{\theta}(1-\hat{\theta})}{N}}$$

Note that σ is the (largest)
when $\hat{\theta} = \frac{1}{2}$, making it
harder to decide whether a
coin is fair, in comparison to
establishing large bias.

More than two outcomes: Beta-Binomial
→ Dirichlet - multinomial

N 'dice' rolls $\{x_1, \dots, x_N\}$ $x_i \in \{1, \dots, k\}$

$$P(\mathcal{D}|\theta) = \prod_k \theta_k^{N_k}$$

$$\sum_k \theta_k = 1$$
$$\sum_k N_k = N$$

For real dice $\theta_1, \dots, \theta_6$

Honestly, if we were dealing with coin tosses or dice rolls, we can have N large enough so that all the Bayesian caution may not be too relevant unless $\theta \sim 1/N$. In practice one may be dealing with text classification or 'words' in DNA. ~~For~~ For words in a natural language $K \sim 10^5$. For DNA sequences 5 bases long, $K = 4^5 \approx 4 \times 10^3$. Some words may be ~~be~~ rare enough not to show up in the training data set, but appear in the test data. We could also have a very common word in the training data be missing in the test set. In all these cases, Bayesian average helps.

prior $\text{Dir}(\theta | \alpha) = \frac{1}{B(\alpha)} \prod_k \theta_k^{\alpha_k - 1}$

Conjugate prior

Posterior $P(\theta | D) \propto P(D | \theta) h(\theta)$

$$\propto \prod_{k=1}^K \theta_k^{\alpha_k + N_k - 1}$$

$$= \text{Dir}(\theta | \alpha_1 + N_1, \dots, \alpha_K + N_K)$$

Optimize subject to constraint $\sum_k \theta_k = 1$

Mode $l(\theta, \lambda) = \sum_k (N_k + \alpha_k - 1) \log \theta_k + \lambda (1 - \sum_k \theta_k)$

$\frac{\partial l}{\partial \theta_k} = 0 \rightarrow \frac{N_k + \alpha_k - 1}{\theta_k} = \lambda \quad \text{or} \quad (N_k + \alpha_k - 1) = \lambda \theta_k$

Sum over k

$$N + \alpha_0 - K = \lambda$$

$$\alpha_0 = \sum_{k=1}^K \alpha_k$$

So $\hat{\theta}_k = \frac{N_k + \alpha_k - 1}{\lambda} = \frac{N_k + \alpha_k - 1}{N + \alpha_0 - K}$

That is the MAP estimate $\hat{\theta}_k = \frac{N_k}{N}$ for MLE.

Posterior predictive

$$p(X=j|D) = \int p(x=j|\theta) p(\theta|D) d\theta$$

$$= \int p(x=j|\theta_j) p(\theta|D) d\theta$$

$$= \int \theta_j p(\theta_j|D) d\theta_j = \bar{\theta}_j$$

$$p(\theta_j|D) = \frac{\int \theta_j^{x_j} \prod_{i \neq j} \theta_i^{x_i} d\theta_j}{\int \prod_{i=1}^K \theta_i^{x_i} d\theta}$$

$$= \frac{\theta_j^{x_j} (1-\theta_j)^{\sum_{i \neq j} x_i + K - 1}}{\int d\theta_j \theta_j^{x_j} (1-\theta_j)^{\sum_{i \neq j} x_i + K - 1}}$$

$$\theta_j^{x_j} (1-\theta_j)^{\sum_{i \neq j} x_i + K - 1}$$

$$x_j = N_j + \alpha_j - 1$$

$$\bar{\theta}_j = \frac{x_j + 1}{x_j + \sum_{i \neq j} x_i + K - 1 + 1} = \frac{N_j + \alpha_j}{N + \alpha_0}$$

Once more, the posterior predictive avoids ~~the~~ the zero count problem [Even when we have a flat prior; $\alpha_j = 1$].

Bag of word model, document classification,
Naive Bayes.

Bag of word model

Mary had 1 little lamb, little lamb, little lamb,
Mary had 1 little lamb, its fleece is white as snow

Drop 'stop words'

Vocab list

mary ₁ lamb ₂ little ₃ big ₄ fleece ₅ white ₆ black ₇ snow ₈

rain ₉ unk. ₁₀ unknown

mary ₁ unk ₁₀ little ₃ lamb ₂ little ₃ lamb ₂ little ₃ lamb ₂

mary ₁ unk ₁₀ little ₃ lamb ₂ unk ₁₀ fleece ₅ white ₆ snow ₈

$$N = 16$$

$$\alpha_j = 1$$

$$\bar{\theta}_j = \frac{N_j + 1}{16 + 10}$$

$$\bar{\theta}_{\text{mary}} = \frac{2+1}{16+10} = \frac{3}{26}$$

$$\bar{\theta}_{\text{big}} = \frac{1}{16+10} = \frac{1}{26}$$

big never shows up

Naive Bayes : $(x_1, \dots, x_D)$ ^{D features}

'Naive' assumption feature conditionally independent (conditioned upon class).

$$P(x|y=c, \theta) = \prod_{j=1}^D P(x_j|y=c, \theta_{j,c})$$

Example: 1) For real-valued ~~features~~ features $\overset{P(x|c, \theta)}{=} \prod_{j=1}^D \mathcal{N}(x_j | \mu_{j,c}, \sigma_{j,c}^2)$

2) Binary features $x_j \in \{0, 1\}$, $P(x|c, \theta) = \prod_{j=1}^D \text{Ber}(x_j | \theta_{j,c})$

Let us see an application to document classification.

Bernoulli document model:

D 'feature' words

$x_1, x_2, \dots, x_{D-1}, x_D$
 $\uparrow \quad \uparrow \quad \quad \quad \uparrow \quad \uparrow$
 word 1 word 2 ... word D-1 word D
 absent present present absent

We throw away the details of word abundance

Bayesian priors

$\pi = (\pi_1, \dots, \pi_C) \rightarrow$ probs of belonging to different class

$p(\pi) \rightarrow \text{Dir}(\alpha)$

$$P(\theta) = \prod_{j=1}^D \prod_{c=1}^C P(\theta_{j,c}) \rightarrow \text{Beta}(\theta_{j,c} | \beta_0, \beta_1)$$

With training data $D \rightarrow N_1, \dots, N_C$
 $\uparrow \qquad \qquad \uparrow$
 numbers
 in different classes

N_{jc} = No. of word j in docs
 of class c

$$P(\theta/D) = P(\pi/D) \prod_{j=1}^D \prod_{c=1}^C P(\theta_{jc}/D)$$

$$P(\pi/D) = \text{Dir}(N_1 + \alpha_1, \dots, N_C + \alpha_C) \propto \bar{\pi}_1^{N_1 + \alpha_1} \dots \bar{\pi}_C^{N_C + \alpha_C}$$

$$P(\theta_{jc}/D) = \text{Beta}(N_c - N_{jc} + \beta_0, N_{jc} + \beta_1) \propto \theta_{jc}^{N_{jc} + \beta_1} (1 - \theta_{jc})^{N_c - N_{jc} + \beta_0}$$

For posterior predictive; integrate out parameters

$$P(y=c|x,D) \propto \int \text{Cat}(y=c|\pi) P(\pi/D) d\pi$$

$$= \prod_{j=1}^D \left[\int \text{Ber}(x_j|y=c, \theta_{jc}) P(\theta_{jc}/D) d\theta_{jc} \right]$$

$$= \prod_{j=1}^D \bar{\theta}_{jc}^{I(x_j=1)} (1 - \bar{\theta}_{jc})^{I(x_j=0)}$$

with $\bar{\pi}_c = \frac{N_c + \alpha_c}{N + \alpha_0}$

$$\bar{\theta}_{jc} = \frac{N_{jc} + \beta_1}{N_c + \beta_0 + \beta_1}$$

$$\alpha_0 = \sum_c \alpha_c$$

Note that $\bar{\theta}_{jc}$ is being used for evaluating the class rather than $(\hat{\theta}_{jc})_{\text{MAP}}$

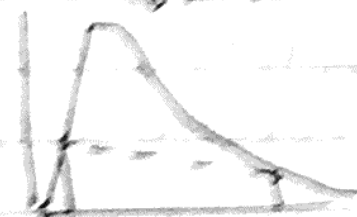
One can use ~~the~~ multinomial model using word frequency, but it could be problematic if certain words are 'bursty': They appear many times when they do in a document, but does not at all appear in most documents.

Back to general issues with Bayesian method. We should take advantage of the fact that Bayesian approach gives us a whole distribution of θ .

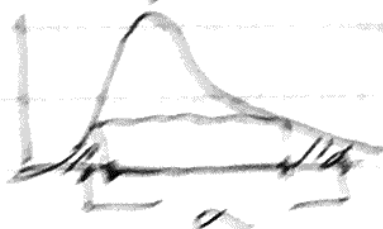
We could do 'confidence intervals' $\rightarrow$ 'credible interval'



If the distn is skewed



Highest posterior density region



If the posterior distribution is not analytically tractable, one can use Monte Carlo to sample from it and get values.

Another important aspect of Bayesian approach is its take on model selection.

Bayesian model Selection

Remember the polynomial fit exercise where choosing different degrees led to models with different complexities. How does Bayesian approach choose among them?

As opposed to the cross validation approach, which requires doing many fits, we compute, for a model m

$$P(m|D) = \frac{P(D|m) p(m)}{\sum_{m'} P(D|m') p(m')}$$

Biggest $P(m|D)$ wins.

For uniform prior ($p(m)$ constant) we just take largest $P(D|m)$

$$P(D|m) = \int P(D|\theta) P(\theta|m) d\theta$$

integrate over
all choices of
parameters

Evidence,
Marginal likelihood,
Integrated likelihood

Bayesian Occam's razor

Occam's razor: One should pick the simplest model that adequately explains the data.

Use $P(D|\hat{\theta}_m)$ to choose between models

Too complex models

$$P(D) = P(y_1) P(y_2|y_1) P(y_3|y_1, y_2) \dots P(y_N|y_1, \dots, y_{N-1})$$

good likelihood by overfitting initial data

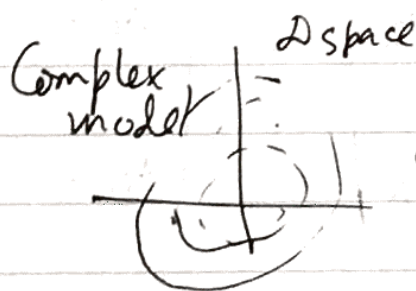
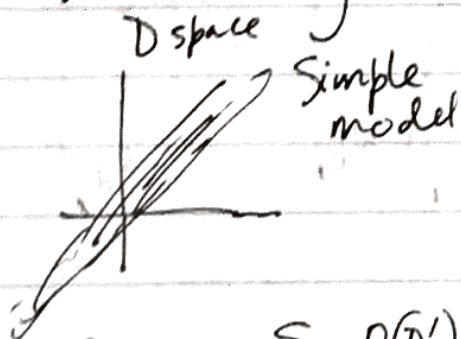
Bad likelihood from later pieces.

Like

training good

validation bad.

Another way is to see it in that



Greater possibility

Since $\sum_{D'} P(D') = 1$

more spread out model usually produces smaller $P(D)$

Example $x_i \sim N(\mu_i, \sigma^2)$

Model 1 $\mu_1 = \mu_2$

$$P(\mu) \propto e^{-\mu^2 / 2\sigma_m^2}$$

Model 2 μ_1, μ_2 independent