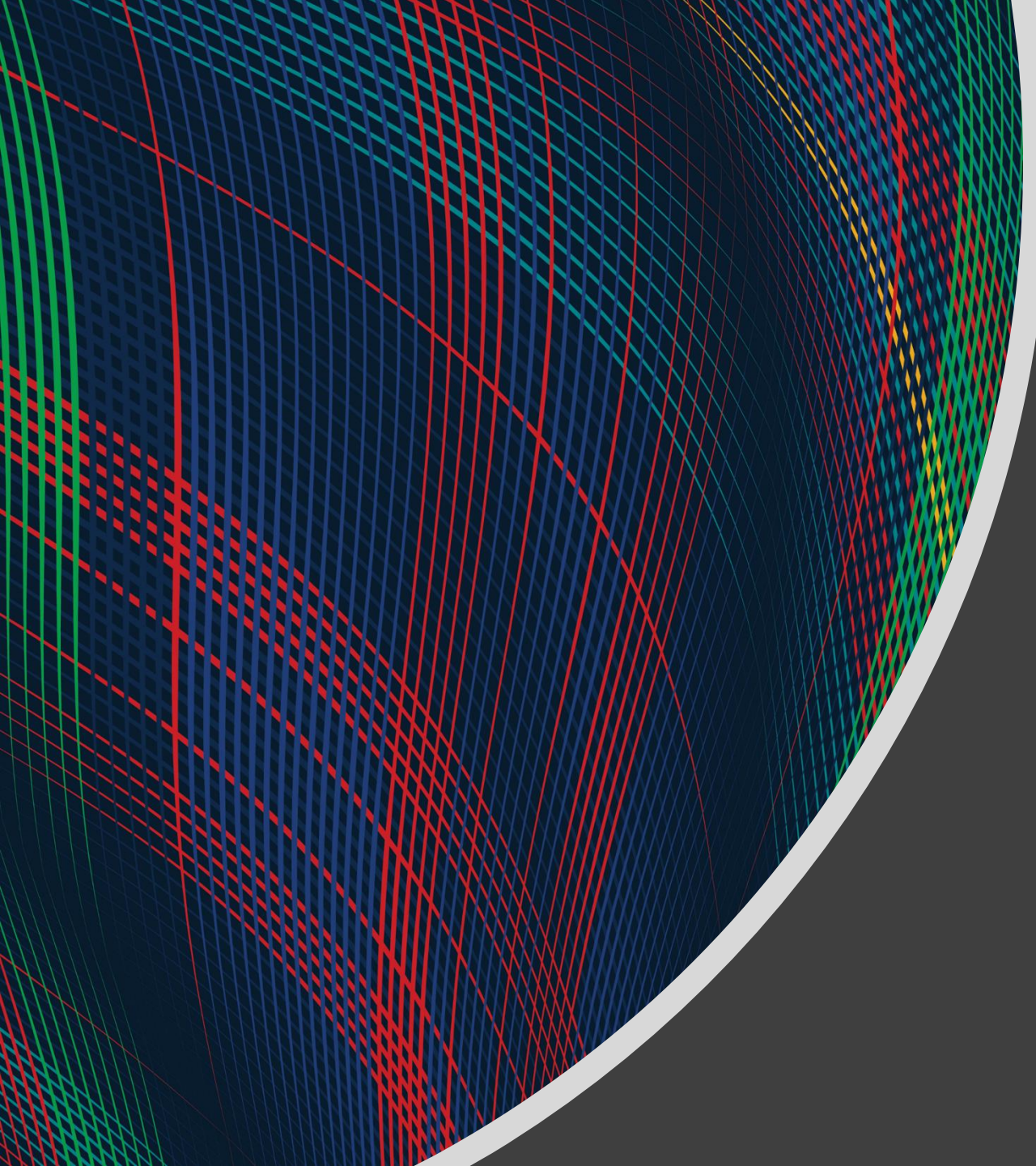




10-315

Introduction to ML

Probabilistic Models:
Generative and Naïve
Bayes

Instructor: Pat Virtue

Poll 1

Let's say we gave an exam that had 3 different versions where the probability of different exams was modeled as a **categorical** random variable Y and the scores from each exam were modeled as $X_{Y=k} \sim \mathcal{N}(\mu_k, \sigma_k^2)$. $\mathcal{D} = \{x^{(i)}, y^{(i)}\}_{i=1}^N$

How many parameters are in this **generative** model?

- A. 2
- B. 3
- ☒ C. 6
- D. 9
- E. 3N

$$\begin{aligned} p(x, y) &= \underbrace{p(x|y)}_{\mathcal{N}(\mu_1, \sigma_1)} p(y) \\ p(y|x) &= \frac{p(x|y)p(y)}{\sum_{y=k}^3 p(x|y)p(y)} \end{aligned}$$

ϕ_1
 ϕ_2
 ϕ_3

Poll 2

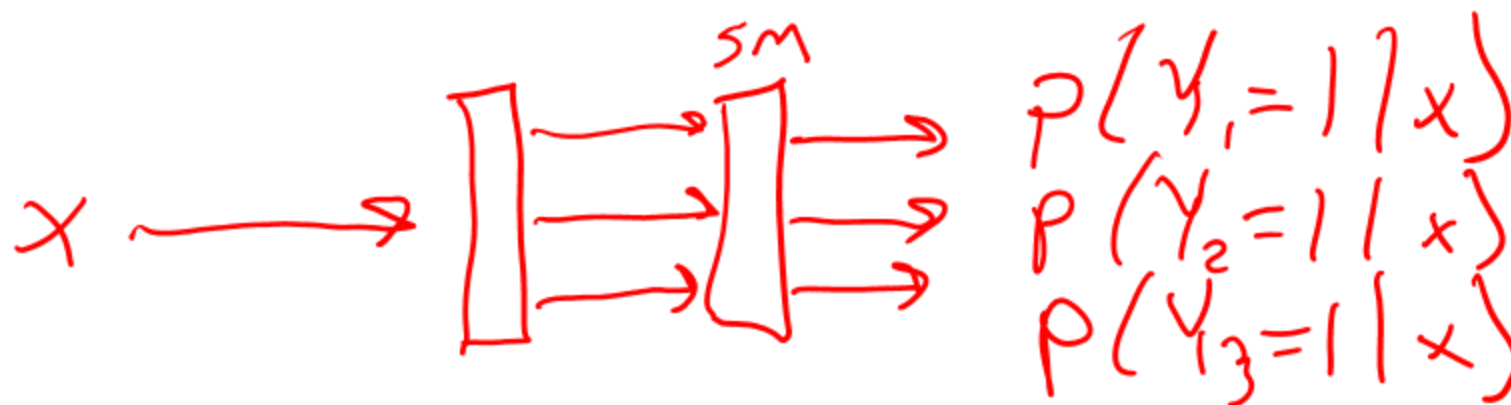
Let's say we gave an exam that had 3 different versions where the probability of different exams Y given the scores from the exam x was modeled as a **multiclass logistic regression** $p(Y_k = 1 \mid x)$ trained on dataset $\mathcal{D} = \{x^{(i)}, y^{(i)}\}_{i=1}^N$

How many parameters are in this **discriminative** model?

- A. 2
- B. 3
- C. 6
- D. 9
- E. $3N$

$x \in \mathbb{R}$ $\rightarrow$ $3 \begin{bmatrix} \square \\ \square \\ \square \end{bmatrix} + 3 \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix}$

$$p(Y_k \mid x) = g_{\text{softmax}}(Wx + b)$$



Poll 3

Which of the following probability models will allow us to sample new data, $\mathcal{D} = \{x_1^{(i)}, x_2^{(i)}, y^{(i)}\}_{i=1}^N$ where x_1, x_2, y are length, weight, and species (cat or dog), respectively? Assume that we have learned all of the corresponding parameters

A) Logistic regression

$$p(Y = cat \mid x_1, x_2, w_1, w_2, b)$$

B) Bernoulli + Gaussian

$$p(Y = cat \mid \phi)$$

$$p(x_1 \mid cat, \mu_{cat,1}, \sigma_{cat,1})$$

$$p(x_2 \mid cat, \mu_{cat,2}, \sigma_{cat,2})$$

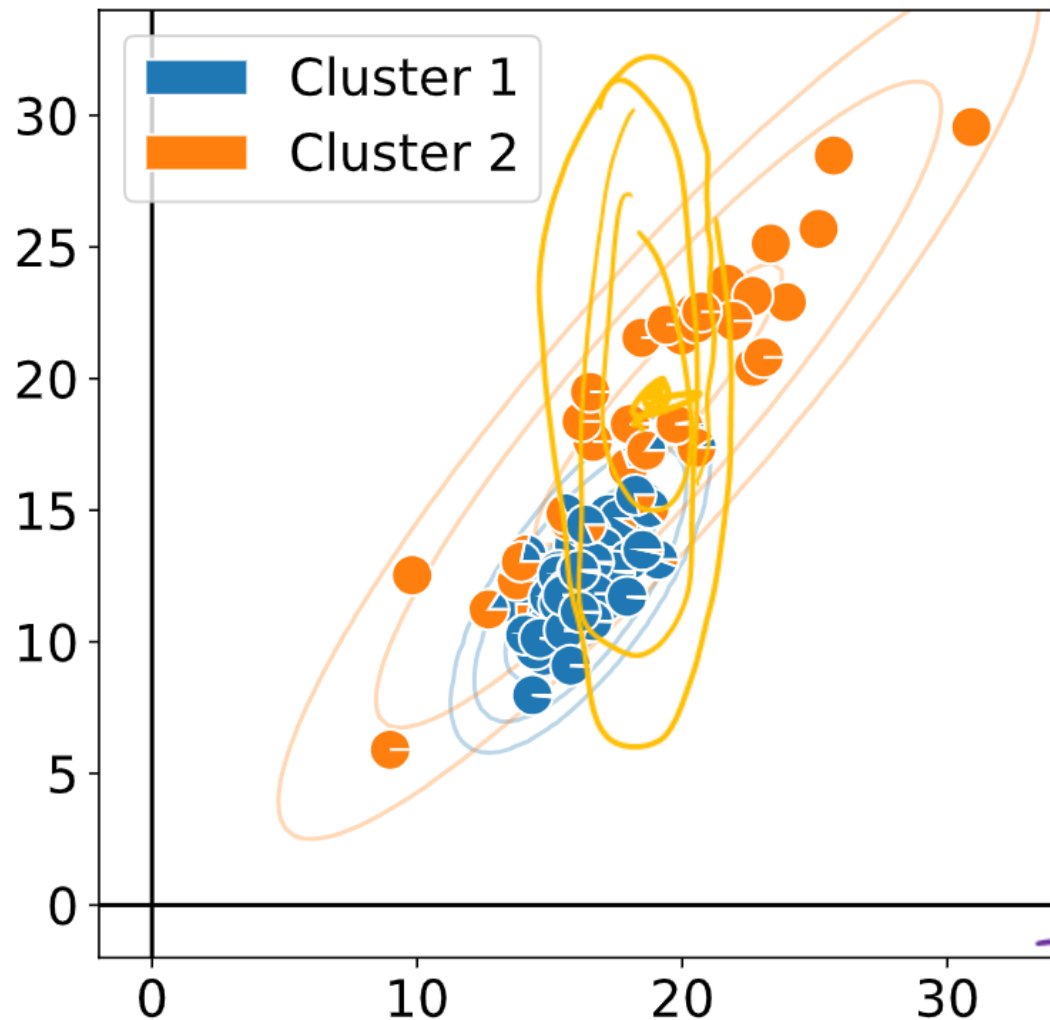
$$p(x_1 \mid dog, \mu_{dog,1}, \sigma_{dog,1})$$

$$p(x_2 \mid dog, \mu_{dog,2}, \sigma_{dog,2})$$

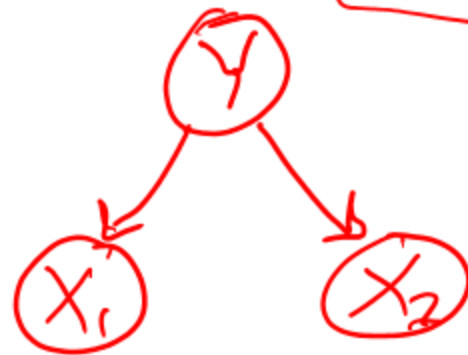
C) None of the above

Sampling Cats and Dogs

$$p(x|y)p(y)$$



$$Y, x_1, x_2$$



$$p(x_1, x_2 | y = \text{dog}) = \mathcal{N} \left(\begin{bmatrix} \mu_1 \\ \mu_2 \end{bmatrix}, \begin{bmatrix} \sigma_1^2 & 0 \\ 0 & \sigma_2^2 \end{bmatrix} \right)$$

B) Bernoulli + Gaussian

Naive Bayes assumption

$$p(Y = \text{cat} | \phi) \leftarrow \text{dog}$$

$$p(x_1 | \text{cat}, \mu_{\text{cat},1}, \sigma_{\text{cat},1})$$

$$p(x_2 | \text{cat}, \mu_{\text{cat},2}, \sigma_{\text{cat},2})$$

$$p(x_1 | \text{dog}, \mu_{\text{dog},1}, \sigma_{\text{dog},1})$$

$$p(x_2 | \text{dog}, \mu_{\text{dog},2}, \sigma_{\text{dog},2})$$

$\leftarrow \mathcal{N}$

Generative Stories

Generative vs Discriminative Modeling

Discriminative: $p(y | x)$ ←

Generative: $p(y | x) = \alpha \boxed{p(x, y)} = \alpha \boxed{p(x | y) p(y)}$

Model assumptions vs Data

- Discriminative: ↓ assump. ↑ data
- Generative: ↑ assumpt. ↓ data

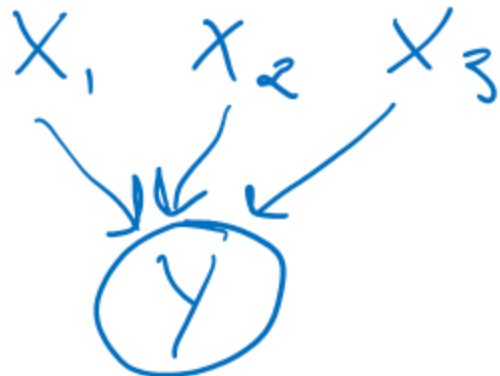
Generative Story Examples

Spam Generation

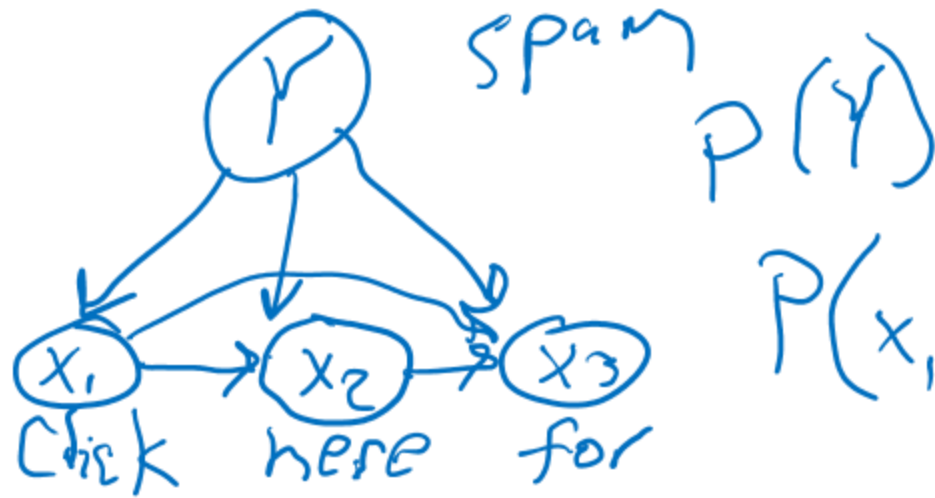
Discriminative

$P(y|x)$

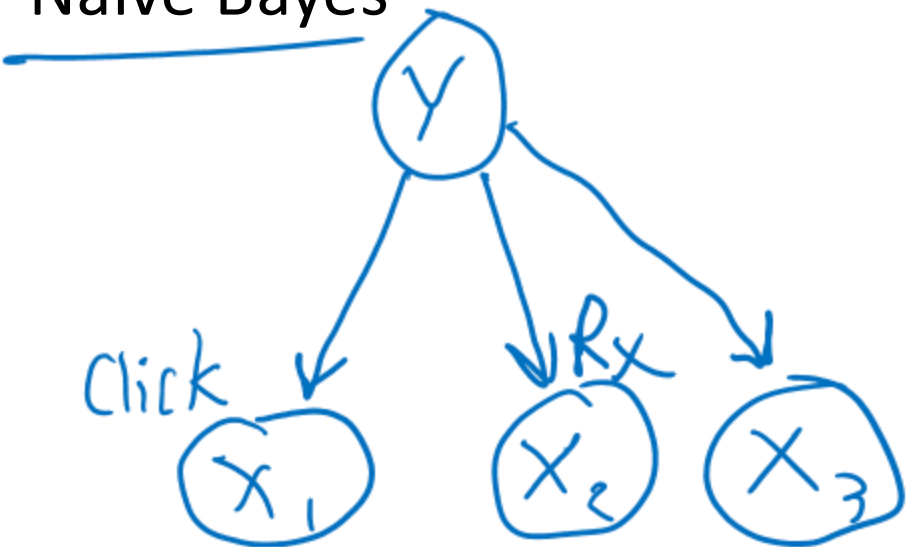




Generative



Generative + Naive Bayes



$P(x_1, x_2, x_3 | y)$

Generative Story Examples

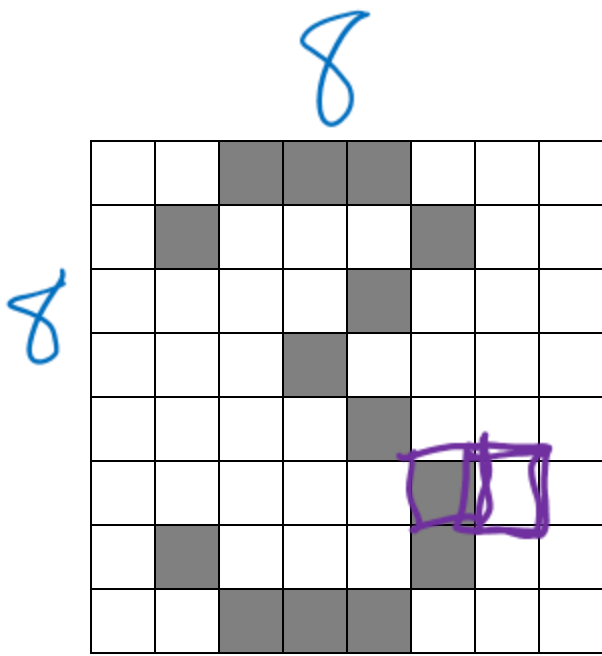
Hand-written digits

y : Digit class: 0-9

x : Pixels in images

$x_1, \dots, x_{64}$

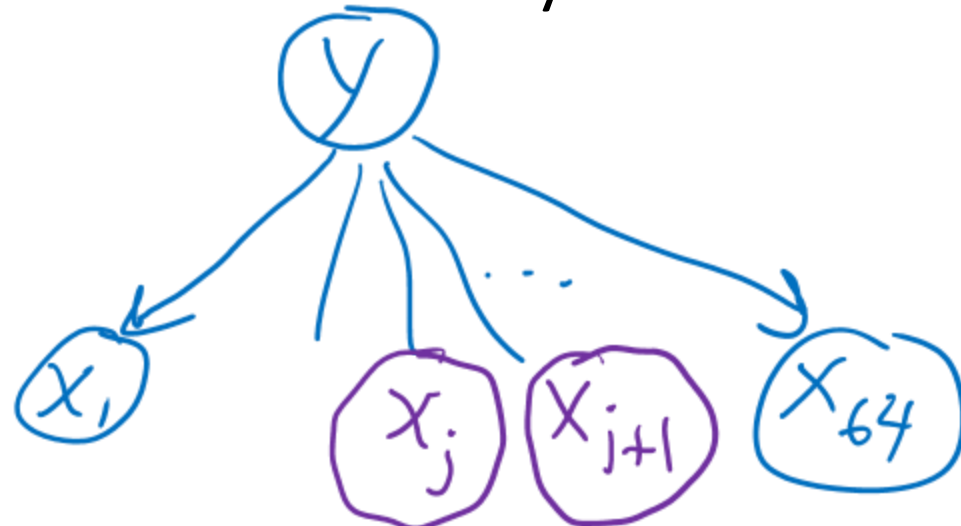
Discriminative



Generative



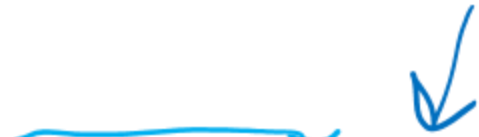
Generative + Naive Bayes



Discriminative and Generative

Two Applications of Bayes Rule

$$p(a | b) = \frac{p(b | a) p(a)}{p(b)}$$

$$p(\theta | \mathcal{D}) = \frac{p(\mathcal{D} | \theta) p(\theta)}{p(\mathcal{D})}$$


$$p(y | x) = \frac{p(x | y) p(y)}{p(x)}$$

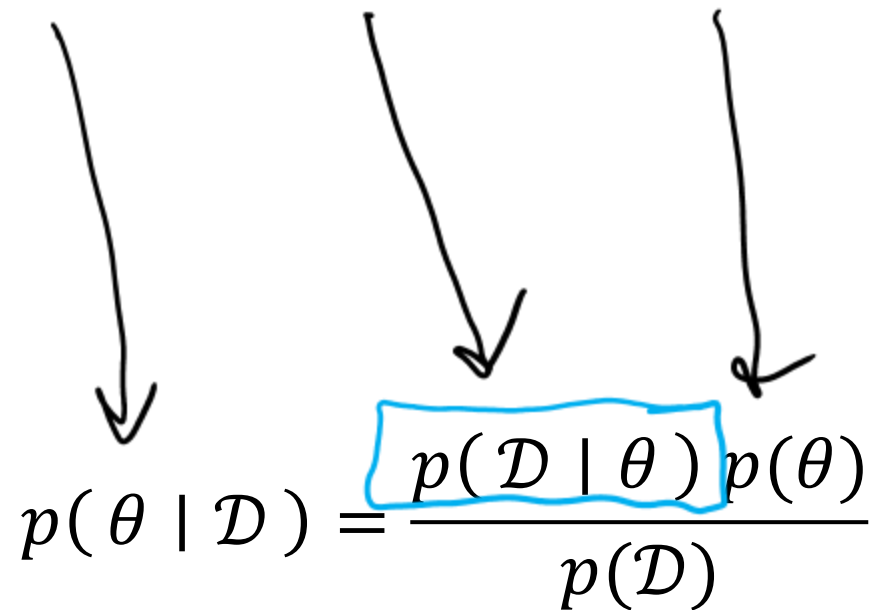

Bayes Rule

Terminology

Posterior

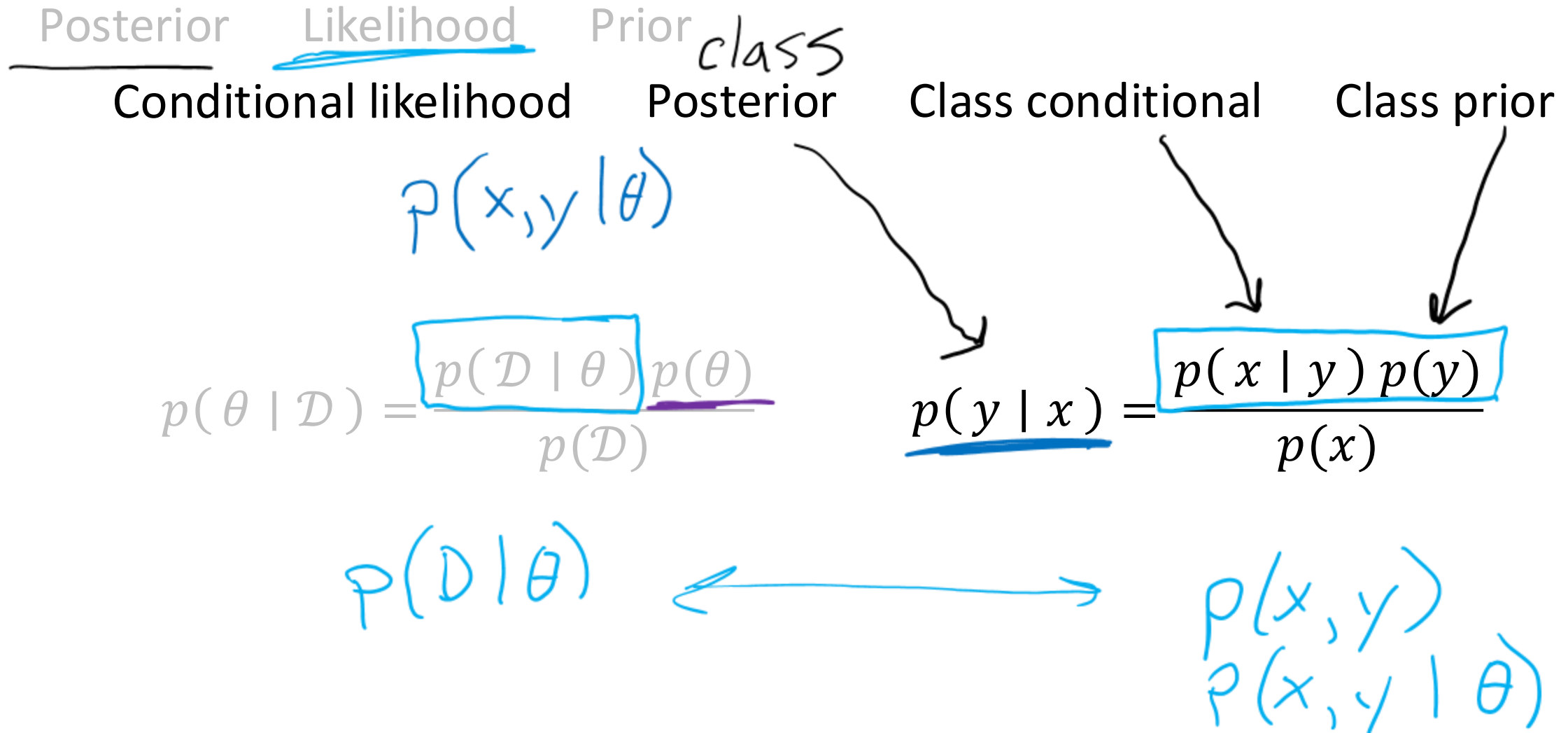
Likelihood

Prior


$$p(\theta | \mathcal{D}) = \frac{p(\mathcal{D} | \theta) p(\theta)}{p(\mathcal{D})}$$

Bayes Rule

Terminology



Bayes Rule

Where did the parameters go?!?

$$p(y | x) = \frac{p(x | y) p(y)}{p(x)}$$

$$p(y | x, \mu_0, \mu, \sigma_0, \sigma, \phi) = \frac{p(x | y, \mu_x, \sigma_y) p(y | \phi)}{p(x)}$$

$$p(y | x, \theta) = \frac{p(x, y | \theta)}{p(x)}$$

$p_\theta(y | x)$

Optimization: Generative vs Discriminative

Discriminative: model $p(y | x, \theta)$ directly

- Learn parameters θ from data

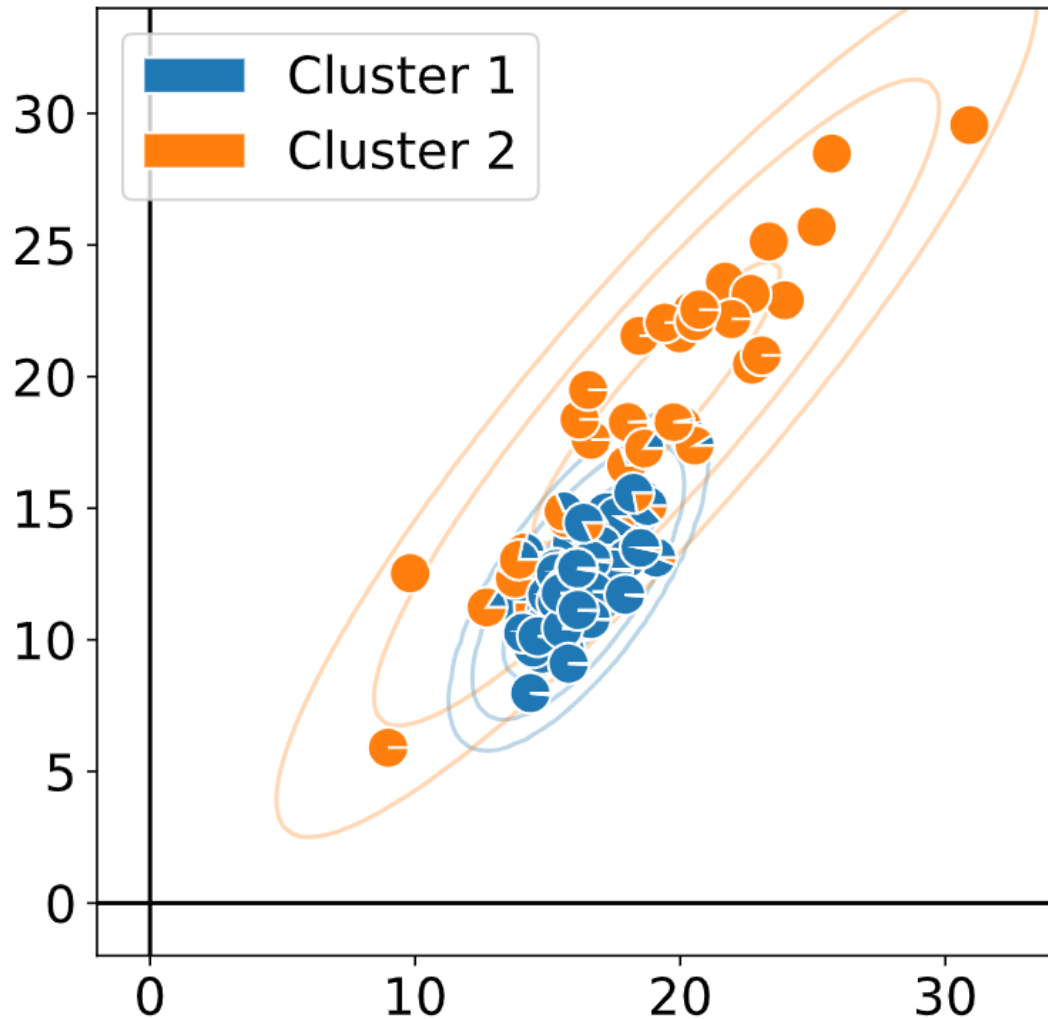
$$J(\theta) \stackrel{L}{=} -\ell \quad \frac{\partial J}{\partial \theta}$$

Generative: model $p(y | \theta_{class})$ and $p(x | y, \theta_{class\ conditional})$

- Learn parameters θ_{class} and $\theta_{class\ conditional}$ from data
- Use Bayes rule to compute $p(y | x, \theta_{class}, \theta_{class\ conditional})$

$$p(y | x, \theta_{class}, \theta_{class\ conditional}) \propto p(x | y, \theta_{class\ conditional}) p(y | \theta_{class})$$

Estimating Cats and Dogs



Bernoulli + Gaussian

$$p(Y = \text{cat} \mid \phi) \quad \phi = \frac{N_{\text{cat}}}{N}$$

$$\rightarrow p(x_1, x_2 \mid Y = \text{cat}, \mu_{\text{cat}}, \Sigma_{\text{cat}})$$

np.mean
np.cov

$$p(x_1, x_2 \mid Y = \text{dog}, \mu_{\text{dog}}, \Sigma_{\text{dog}})$$

Reminder: Likelihood

Simple example with three i.i.d. samples of $Y \sim \text{Bernoulli}(\phi)$

$$\mathcal{D} = \{\underline{y^{(i)}}\}_{i=1}^3$$

Likelihood

$$p(\mathcal{D}; \theta)$$

$$= p(y^{(1)}, y^{(2)}, y^{(3)}; \phi^{(1)}, \phi^{(2)}, \phi^{(3)})$$
 Expand notation

$$= p(y^{(1)}, y^{(2)}, y^{(3)}; \phi)$$
 Identically distributed

$$= p(y^{(1)}; \phi) p(y^{(2)}; \phi) p(y^{(3)}; \phi)$$
 Independent

$$= \prod_{i=1}^3 p(y^{(i)}; \phi)$$

Generative MLE

$$L(\phi, \Theta) = p(\mathcal{D} \mid \phi, \Theta)$$

$$= \prod_{n=1}^N p(\mathcal{D}^{(n)} \mid \phi, \Theta) \quad \text{i.i.d assumption}$$

$$= \prod_{n=1}^N p(\underline{y^{(n)}, \mathbf{x}^{(n)}} \mid \phi, \Theta)$$

Generative model

$$= \prod_{n=1}^N p(\underline{y^{(n)} \mid \phi}) p(\underline{\mathbf{x}^{(n)} \mid y^{(n)}, \Theta})$$

Generative model story

$$= \prod_{n=1}^N p(\underline{y^{(n)} \mid \phi}) p(\underline{x_1^{(n)}, x_2^{(n)}, \dots, x_M^{(n)} \mid y^{(n)}, \Theta})$$

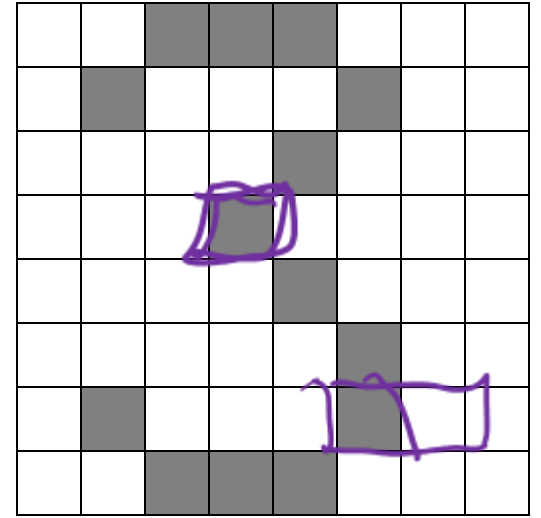
$$\begin{aligned} \mathcal{D} &= \{y^{(i)}, x^{(i)}\}_{i=1}^N \\ y^{(i)} &\in \{0,1\} \\ \mathbf{x}^{(i)} &\in \{0,1\}^M \\ \phi &\in [0,1] \\ \Theta &\in [0,1]^{M \times 2} \end{aligned}$$

Naïve Bayes

Multivariate Generative Models

Hand-written digits: How many parameters?

- $P(Y)$
- $P(\mathbf{X} \mid Y = 3)$
 $= p(X_1, X_2, \dots, X_{64} \mid Y = 3)$



Naïve Bayes assumption (bag of pixels)

- $P(Y)$
- $P(\mathbf{X} \mid Y = 3)$
 $= p(X_1 \mid Y = 3) p(X_2 \mid Y = 3) \dots p(X_{64} \mid Y = 3)$

Conditional Independence and Naïve Bayes

Independence

$$A \perp\!\!\!\perp B$$

$$P(A | B) = P(A)$$

$$P(B | A) = P(B)$$

$$P(A, B) = P(A)P(B)$$

Conditional independence

$$A \perp\!\!\!\perp B \mid C$$

$$P(A | B, \underline{C}) = P(A | \underline{C})$$

$$P(B | A, \underline{C}) = P(B | \underline{C})$$

$$P(A, B | \underline{C}) = P(A | \underline{C})P(B | \underline{C})$$

Naïve Bayes assumption

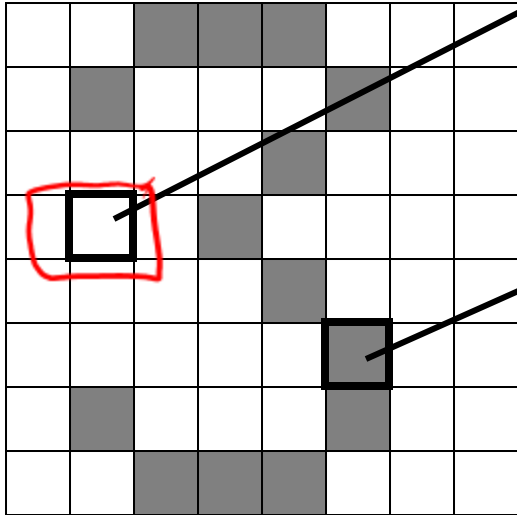
$$\forall j \neq k$$

$$x_i \perp\!\!\!\perp x_j \mid y$$

$$P(x_1 x_2 x_3 \dots x_n | y) = \prod_{j=1}^n P(x_j | y)$$

Naïve Bayes for Digits

y	$P(Y)$
1	0.1
2	0.1
3	0.1
4	0.1
5	0.1
6	0.1
7	0.1
8	0.1
9	0.1
0	0.1



y	$P(X_{3,1} = 1 y)$
1	0.01
2	0.05
3	0.05
4	0.30
5	0.80
6	0.90
7	0.05
8	0.60
9	0.50
0	0.80

y	$P(X_{5,5} = 1 y)$
1	0.05
2	0.01
3	0.90
4	0.80
5	0.90
6	0.90
7	0.25
8	0.85
9	0.60
0	0.80

$$\phi_{y=3, x_{3,1}}$$



SPAM Classification

Recitation Exercise

$P(Y = 0, X_1, \dots, X_M)$	$P(Y = 1, X_1, \dots, X_M)$

$P(Y = 0 \mid X_1, \dots, X_M)$	$P(Y = 1 \mid X_1, \dots, X_M)$

Naïve Bayes MLE

$$L(\phi, \Theta) = p(\mathcal{D} \mid \phi, \Theta)$$

$$= \prod_{n=1}^N p(\mathcal{D}^{(n)} \mid \phi, \Theta) \quad \text{i.i.d assumption}$$

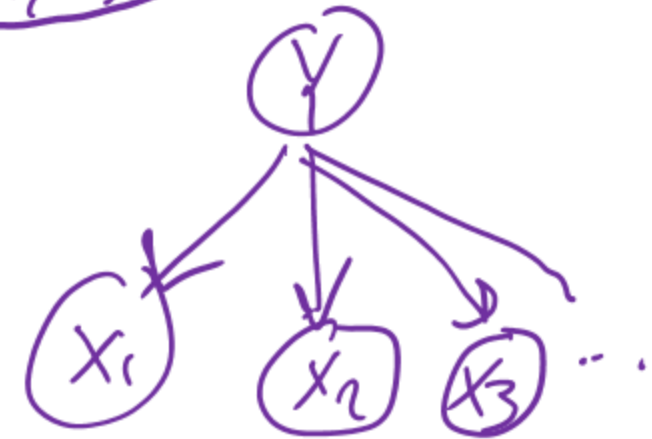
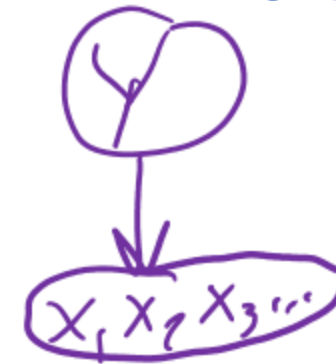
$$= \prod_{n=1}^N p(y^{(n)}, \mathbf{x}^{(n)} \mid \phi, \Theta)$$

$$= \prod_{n=1}^N p(y^{(n)} \mid \phi) p(\mathbf{x}^{(n)} \mid y^{(n)}, \Theta) \quad \text{Generative model}$$

$$= \prod_{n=1}^N p(y^{(n)} \mid \phi) \underbrace{p(x_1^{(n)}, x_2^{(n)}, \dots, x_M^{(n)} \mid y^{(n)}, \Theta)}$$

$$= \prod_{n=1}^N p(y^{(n)} \mid \phi) \prod_{m=1}^M p(\underbrace{x_m^{(n)}}_{\text{feature}} \mid y^{(n)}, \theta_{m,y}) \quad \text{Naïve Bayes}$$

$$\begin{aligned} \mathcal{D} &= \{y^{(n)}, \mathbf{x}^{(n)}\}_{n=1}^N \\ y^{(n)} &\in \{0,1\} \\ \mathbf{x}^{(n)} &\in \{0,1\}^M \\ \phi &\in [0,1] \\ \Theta &\in [0,1]^{M \times 2} \end{aligned}$$



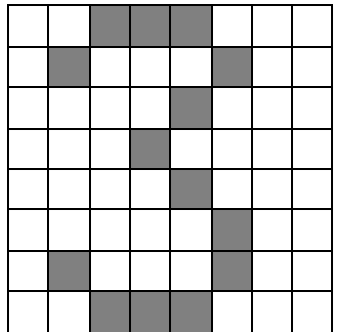
Generative Models

SPAM:

- Class distribution: $Y \sim \text{Bern}(\phi)$
- Class conditional distribution: $X_m \sim \text{Bern}(\theta_{m,y})$
- Naïve Bayes X_i conditionally independent X_j given Y for all $i \neq j$
$$p(X_i, X_j | Y) = p(X_i | Y) | p(X_j | Y)$$

Digits:

- Class distribution: $Y \sim \text{Categorical}(\boldsymbol{\phi})$
- Class conditional distribution: $X_m \sim \text{Bern}(\theta_{m,y})$
- Naïve Bayes X_i conditionally independent X_j given Y for all $i \neq j$
$$p(X_i, X_j | Y) = p(X_i | Y) | p(X_j | Y)$$



Generative Models with Continuous Features

Iris dataset:

- Class distribution: $Y \sim \text{Bern}(\phi)$
- Class conditional distribution: Multivariate Gaussian $\mathbf{X} \sim \mathcal{N}(\boldsymbol{\mu}_y, \boldsymbol{\Sigma}_y)$
- Naïve Bayes assumption?

Categorical

x_1
 x_2
 x_3
 x_4

$$\begin{bmatrix} \sigma_1^2 & 0 & 0 & 0 \\ 0 & \sigma_2^2 & 0 & 0 \\ 0 & 0 & \sigma_3^2 & 0 \\ 0 & 0 & 0 & \sigma_4^2 \end{bmatrix}$$

Poll 5

Iris dataset:

- Class distribution: $Y \sim \text{Bern}(\phi)$
- Class conditional distribution: $\mathbf{X} \sim \mathcal{N}(\boldsymbol{\mu}_y, \boldsymbol{\Sigma}_y)$
- Naïve Bayes assumption?

Which of the following pairs of Gaussian class conditional distributions satisfy the Naïve Bayes assumptions? Select ALL that apply.

$$A. \quad \boldsymbol{\mu}_{y=0} = \begin{bmatrix} -1 \\ 0 \end{bmatrix}, \boldsymbol{\Sigma}_{y=0} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}, \quad \boldsymbol{\mu}_{y=1} = \begin{bmatrix} 1 \\ 0 \end{bmatrix}, \boldsymbol{\Sigma}_{y=1} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$$

$$B. \quad \boldsymbol{\mu}_{y=0} = \begin{bmatrix} -1 \\ 0 \end{bmatrix}, \boldsymbol{\Sigma}_{y=0} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}, \quad \boldsymbol{\mu}_{y=1} = \begin{bmatrix} 1 \\ 0 \end{bmatrix}, \boldsymbol{\Sigma}_{y=1} = \begin{bmatrix} 3 & 0 \\ 0 & 3 \end{bmatrix}$$

$$C. \quad \boldsymbol{\mu}_{y=0} = \begin{bmatrix} -1 \\ 0 \end{bmatrix}, \boldsymbol{\Sigma}_{y=0} = \begin{bmatrix} 1 & -1 \\ 1 & 2 \end{bmatrix}, \quad \boldsymbol{\mu}_{y=1} = \begin{bmatrix} 1 \\ 0 \end{bmatrix}, \boldsymbol{\Sigma}_{y=1} = \begin{bmatrix} 1 & -1 \\ 1 & 2 \end{bmatrix}$$

$$D. \quad \boldsymbol{\mu}_{y=0} = \begin{bmatrix} -1 \\ 0 \end{bmatrix}, \boldsymbol{\Sigma}_{y=0} = \begin{bmatrix} 1 & -1 \\ 1 & 2 \end{bmatrix}, \quad \boldsymbol{\mu}_{y=1} = \begin{bmatrix} 1 \\ 0 \end{bmatrix}, \boldsymbol{\Sigma}_{y=1} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$$

Class-conditional Gaussian Distributions

Iris dataset:

- Class distribution: $Y \sim \text{Bern}(\phi)$ (or Categorical)
- Class conditional distribution: $\mathbf{X} \sim \mathcal{N}(\boldsymbol{\mu}_y, \boldsymbol{\Sigma}_y)$

- Naïve Bayes assumption: Σ are all diagonal

- Linear Decision Boundary: $\Sigma_{Y=0} = \Sigma_{Y=1}$

- Quadratic Decision Boundary: $\Sigma_{Y=0} \neq \Sigma_{Y=1}$

Generative + MAP

MLE vs MAP vs Generative vs Discriminative

Maximum likelihood estimation

Maximum *a posteriori* estimation

Discriminative

Models only the
conditional
likelihood,
 $\prod p(y | x, \theta)$

$$p(\mathcal{D}^* | \theta) \quad \text{* conditional likelihood}$$

$$= \prod_{i=1}^N p(y^{(i)} | x^{(i)}, \theta)$$

$$p(\theta) p(\mathcal{D}^* | \theta) \quad \text{* conditional likelihood}$$

$$= p(\theta) \prod_{i=1}^N p(y^{(i)} | x^{(i)}, \theta)$$

Generative

Models the full
joint likelihood,
 $\prod p(x, y | \theta)$

$$p(\mathcal{D} | \theta) \quad \text{* actual likelihood}$$

$$= \prod_{i=1}^N p(x^{(i)}, y^{(i)} | \theta)$$

$$= \prod_{i=1}^N \underbrace{p(x^{(i)} | y^{(i)}, \theta)}_{\text{class-conditional}} \underbrace{p(y^{(i)} | \theta)}_{\text{class prior}}$$

$$p(\theta) p(\mathcal{D} | \theta) \quad \text{* actual likelihood}$$

$$= p(\theta) \prod_{i=1}^N p(x^{(i)}, y^{(i)} | \theta)$$

$$= p(\theta) \prod_{i=1}^N \underbrace{p(x^{(i)} | y^{(i)}, \theta)}_{\text{class-conditional}} \underbrace{p(y^{(i)} | \theta)}_{\text{class prior}}$$

Reminder: Prior Distributions for MAP

If the prior $p(\theta)$ is uniform, then MLE and MAP are the same!

$$p(\mathcal{D} \mid \phi) p(\phi) = \prod_i^N p(y^{(i)} \mid \phi) p(\phi) = \phi^{N_{y=1}} (1 - \phi)^{N_{y=0}} p(\phi)$$

Conjugate priors: when the prior and the posterior distributions are in the same family

Bernoulli likelihood with a Beta prior has Beta posterior

Categorical likelihood with a Dirichlet prior has Dirichlet posterior

Gaussian likelihood with a Gaussian prior has Gaussian posterior

<https://www.desmos.com/calculator/kr7m2m6cf7>