

Hi

Announcements

Cody homework emailed.

Due Monday and Wednesday before class

Email me, if just attending as that way you can look at homework.

Piazza to come

Podcast might work eventually.

Introduction to ML, lecture 20 ;-)

Grade last year (A+ 19, A 20, A- 13, B+ 7, S 1, W 1)

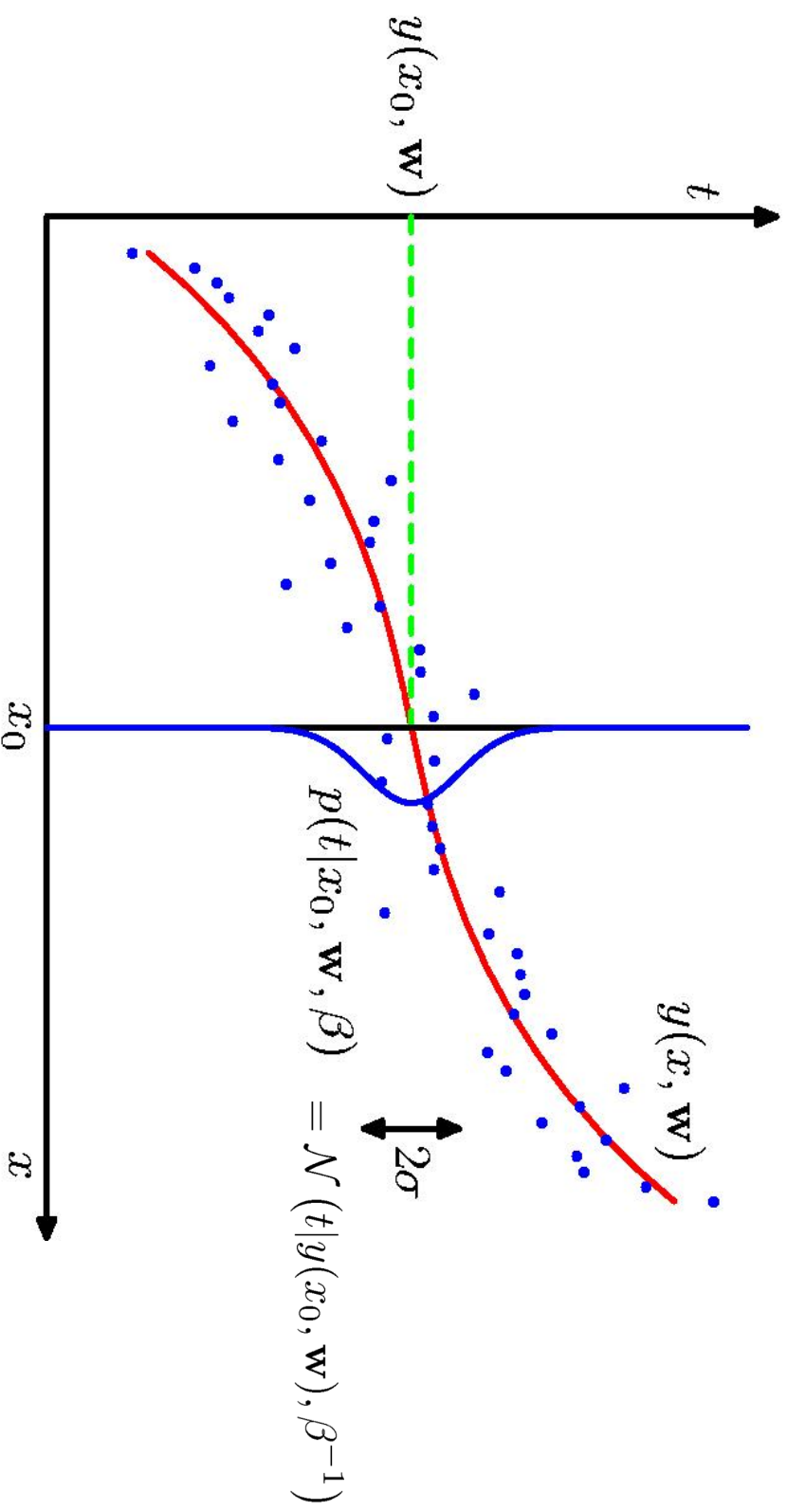
Today:

- Gaussian 1.2
- Decision theory 1.5
- Information theory 1.6
- Homework
- Gaussian 2.3

Monday

Gaussian 2.3, Non parametric 1.5, Linear models for regression 3

Curve Fitting Re-visited, Bishop1.2.5



Maximum Likelihood Bishop 1.2.5

- Model
$$t = w^T x + n \quad n \sim \mathcal{N}(0, \sigma^2) = \mathcal{N}(0, \beta^{-1})$$

$$t \sim \mathcal{N}(w^T x, \sigma^2)$$

- Likelihood
$$L(w) = \prod_{n=1}^N \mathcal{N}(t_n | w^T x_n, \sigma^2) = \prod_{n=1}^N \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{1}{2}\left(\frac{t_n - w^T x_n}{\sigma}\right)^2}$$

$$L(w) = -\log L(w) = \frac{1}{2\sigma^2} \sum (t_n - w^T x_n)^2 + \frac{N}{2} \log \sigma^2 + \frac{N}{2} \log 2\pi$$

- differentiation

$$\frac{\partial L}{\partial \sigma^2} = \frac{1}{2\sigma^4} \sum (t_n - w^T x_n)^2 + \frac{N}{2} \frac{1}{\sigma^2} = 0$$

$$\sigma^2 = \frac{\sum (t_n - w^T x_n)^2}{N}$$

$$\frac{dL}{dw} = \frac{1}{\sigma^2} \sum (t_n - w^T x_n) x_n$$

Gradient Minimization

Maximum Likelihood

$$p(\mathbf{t}|\mathbf{x}, \mathbf{w}, \beta) = \prod_{n=1}^N \mathcal{N}(t_n | y(x_n, \mathbf{w}), \beta^{-1}). \quad (1.61)$$

As we did in the case of the simple Gaussian distribution earlier, it is convenient to maximize the logarithm of the likelihood function. Substituting for the form of the Gaussian distribution, given by (1.46), we obtain the log likelihood function in the form

$$\ln p(\mathbf{t}|\mathbf{x}, \mathbf{w}, \beta) = -\frac{\beta}{2} \sum_{n=1}^N \{y(x_n, \mathbf{w}) - t_n\}^2 + \frac{N}{2} \ln \beta - \frac{N}{2} \ln(2\pi). \quad (1.62)$$

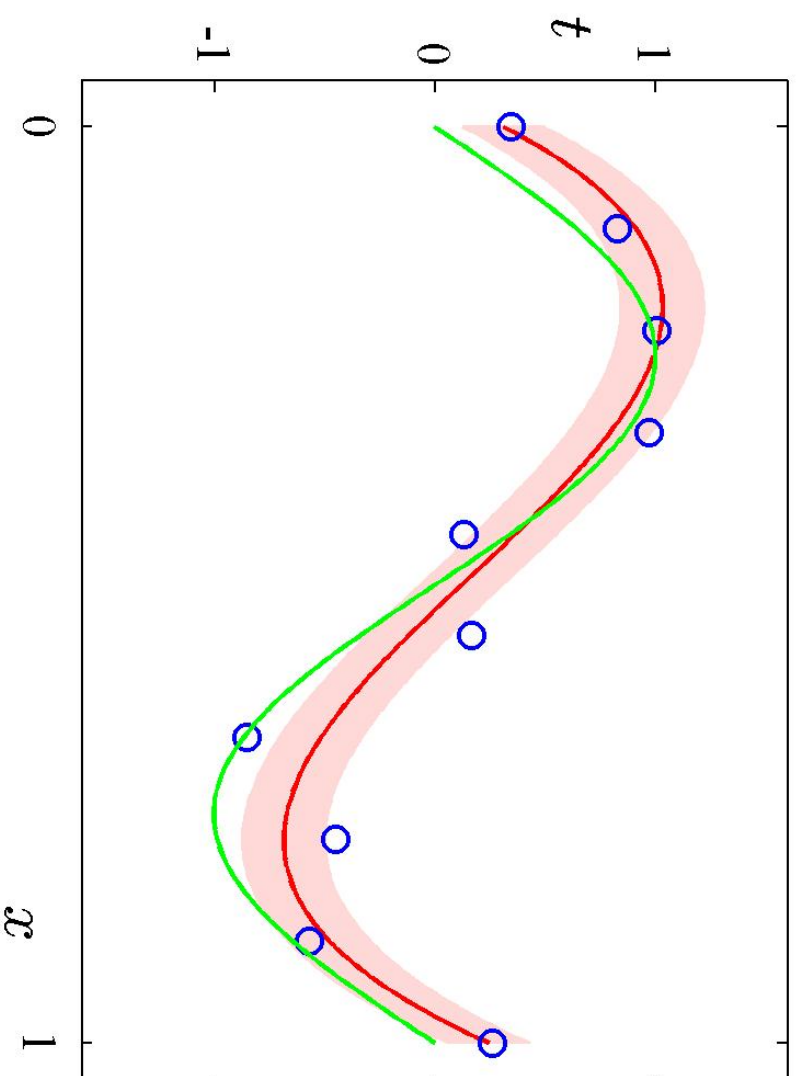
$$\underbrace{\sum}_{\text{ML}} \frac{1}{\beta_{\text{ML}}} = \frac{1}{N} \sum_{n=1}^N \{y(x_n, \mathbf{w}_{\text{ML}}) - t_n\}^2. \quad (1.63)$$

Giving estimates of $\mathbf{w}$ and beta, we can predict

$$p(t|x, \mathbf{w}_{\text{ML}}, \beta_{\text{ML}}) = \mathcal{N}(t | y(x, \mathbf{w}_{\text{ML}}), \beta_{\text{ML}}^{-1}). \quad (1.64)$$

Predictive Distribution

$$p(t|x, \mathbf{w}_{\text{ML}}, \beta_{\text{ML}}) = \mathcal{N}(t|y(x, \mathbf{w}_{\text{ML}}), \beta_{\text{ML}}^{-1})$$



MAP: A Step towards Bayes 1.2.5

$$p(\mathbf{w}|\alpha) = \mathcal{N}(\mathbf{w}|\mathbf{0}, \alpha^{-1}\mathbf{I}) = \left(\frac{\alpha}{2\pi}\right)^{(M+1)/2} \exp\left\{-\frac{\alpha}{2}\mathbf{w}^T\mathbf{w}\right\}$$

$$p(\mathbf{w}|\mathbf{x}, \mathbf{t}, \alpha, \beta) \propto p(\mathbf{t}|\mathbf{x}, \mathbf{w}, \beta)p(\mathbf{w}|\alpha) = \prod_{n=1}^N \frac{\sqrt{\beta}}{\sqrt{2\pi}} e^{-\beta(\mathbf{w}^T\mathbf{x}_n - t_n)^2} \cdot \frac{\sqrt{\alpha}}{\sqrt{2\pi}} e^{-\frac{\alpha}{2}\mathbf{w}^T\mathbf{w}}$$

$\propto \frac{1}{2} \sum_{n=1}^N (\mathbf{w}^T\mathbf{x}_n - t_n)^2 + \frac{\alpha}{2} \mathbf{w}^T\mathbf{w}$

$$\beta \tilde{E}(\mathbf{w}) = \frac{\beta}{2} \sum_{n=1}^N \{y(x_n, \mathbf{w}) - t_n\}^2 + \frac{\alpha}{2} \mathbf{w}^T\mathbf{w}$$

Determine $\mathbf{w}_{\text{MAP}}$ by minimizing regularized sum-of-squares error, $\tilde{E}(\mathbf{w})$.

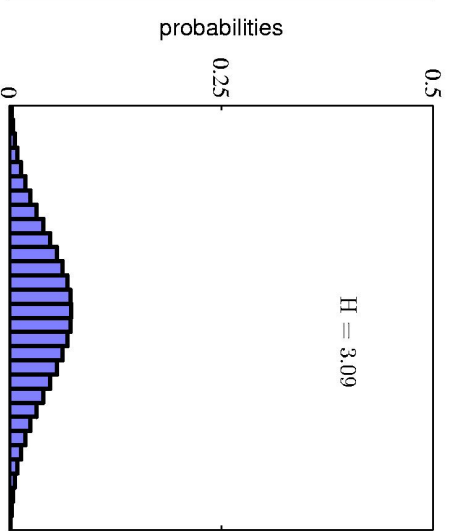
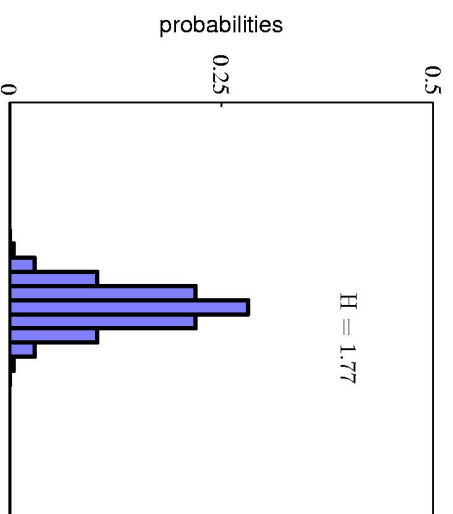
Regularized sum of squares

Entropy 1.6

$$H[x] = - \sum_x \bar{p}(x) \log_2 p(x)$$

Important quantity in

- coding theory
- statistical physics
- machine learning



Differential Entropy

Put bins of width Δ along the real line

$$\lim_{\Delta \rightarrow 0} \left\{ - \sum_i p(x_i) \Delta \ln p(x_i) \right\} = - \int p(x) \ln p(x) dx$$

For fixed σ^2 differential entropy maximized when

in which case

$$p(x) = \mathcal{N}(x|\mu, \sigma^2)$$

$$H[x] = \frac{1}{2} \{1 + \ln(2\pi\sigma^2)\}.$$

The Kullback-Leibler Divergence

P true distribution, q is approximating distribution

$$\begin{aligned}\text{KL}(p\|q) &= - \int p(\mathbf{x}) \ln q(\mathbf{x}) \, d\mathbf{x} - \left(- \int p(\mathbf{x}) \ln \underline{p(\mathbf{x})} \, d\mathbf{x} \right) \\ &= \underbrace{- \int p(\mathbf{x}) \ln \left\{ \frac{q(\mathbf{x})}{p(\mathbf{x})} \right\} \, d\mathbf{x}}\end{aligned}$$

$$\text{KL}(p\|q) \simeq \frac{1}{N} \sum_{n=1}^N \{ -\ln q(\mathbf{x}_n | \boldsymbol{\theta}) + \ln p(\mathbf{x}_n) \}$$

$$\text{KL}(p\|q) \geq 0 \qquad \text{KL}(p\|q) \neq \text{KL}(q\|p)$$

Decision Theory

1.5

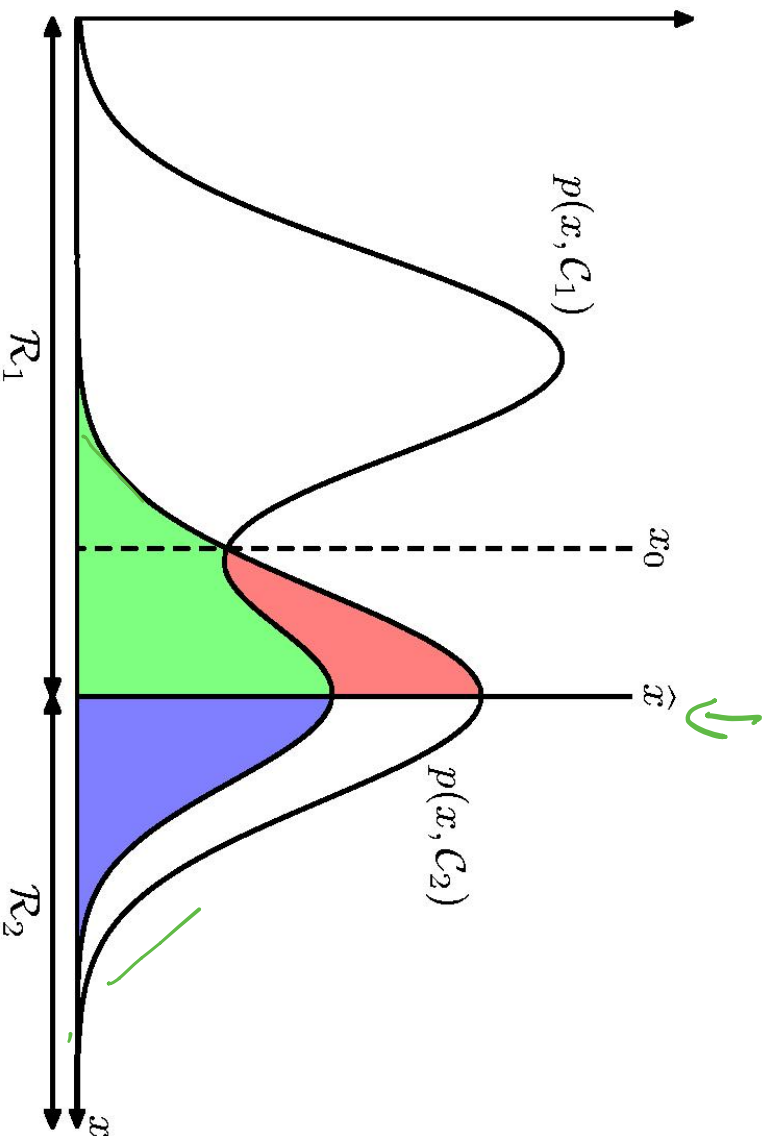
Inference step

Determine either $p(t|\mathbf{x})$ or $p(\mathbf{x}, t)$.

Decision step

For given $\mathbf{x}$, determine optimal t .

Minimum Misclassification Rate

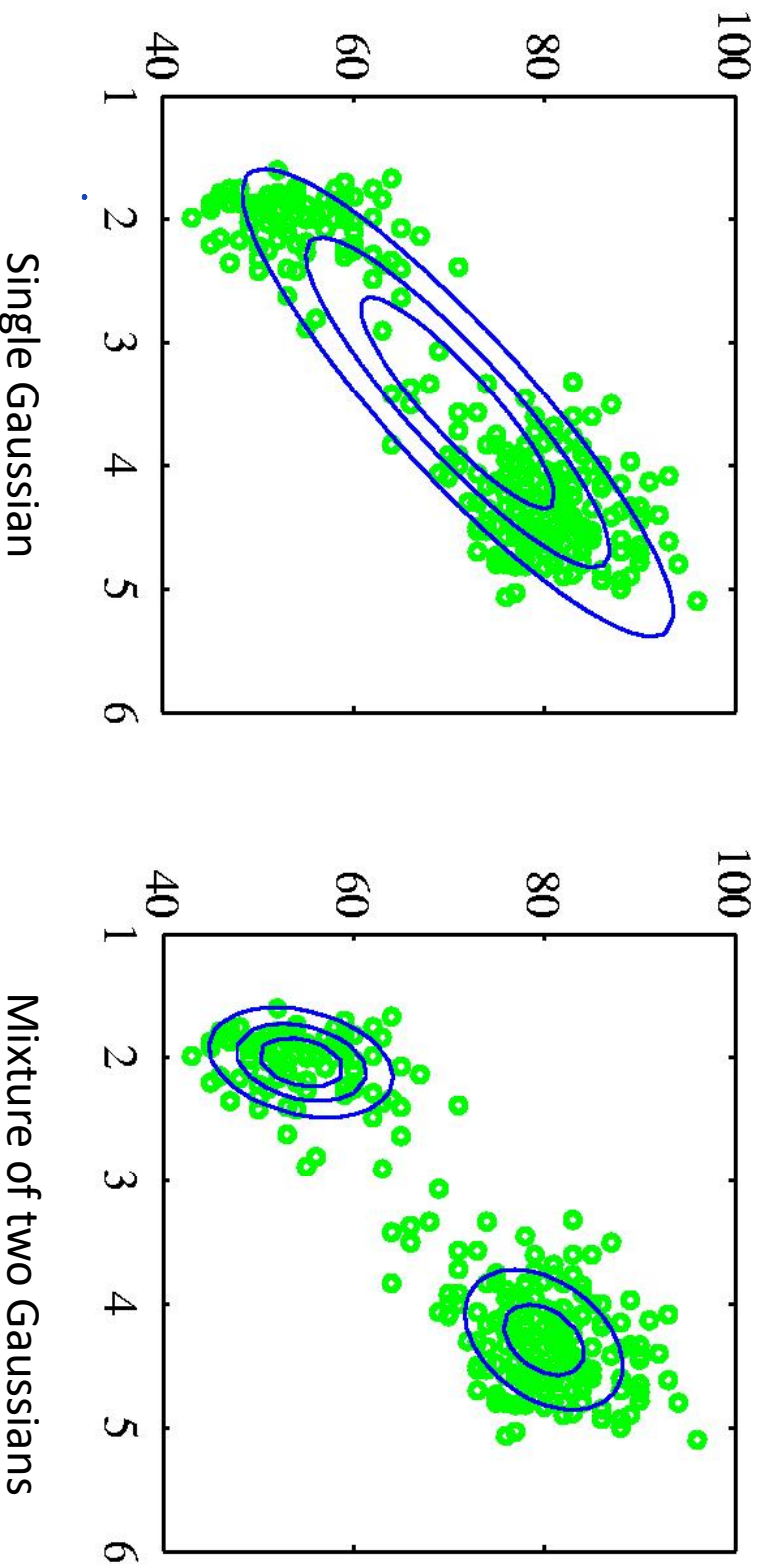


$$\begin{aligned}
 p(\text{mistake}) &= p(\mathbf{x} \in \mathcal{R}_1, \mathcal{C}_2) + p(\mathbf{x} \in \mathcal{R}_2, \mathcal{C}_1) \\
 &= \int_{\mathcal{R}_1} p(\mathbf{x}, \mathcal{C}_2) d\mathbf{x} + \int_{\mathcal{R}_2} p(\mathbf{x}, \mathcal{C}_1) d\mathbf{x}.
 \end{aligned}$$

Mixtures of Gaussians (Bishop 2.3.9)

Old Faithful geyser:

The time between eruptions has a bimodal distribution, with the mean interval being either 65 or 91 minutes, and is dependent on the length of the prior eruption. Within a margin of error of ± 10 minutes, Old Faithful will erupt either 65 minutes after an eruption lasting less than $2\frac{1}{2}$ minutes, or 91 minutes after an eruption lasting more than $2\frac{1}{2}$ minutes.



Mixtures of Gaussians (Bishop 2.3.9)

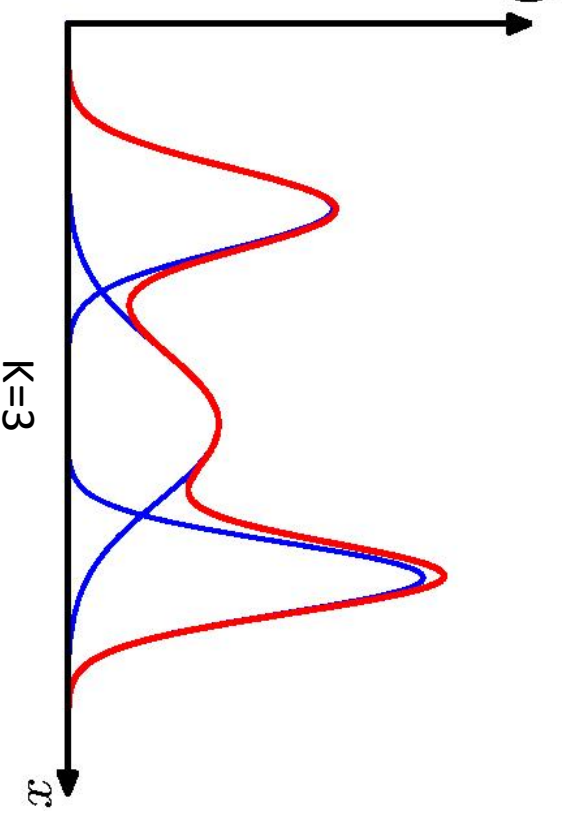
- Combine simple models $p(x)$ into a complex model:

$$p(\mathbf{x}) = \sum_{k=1}^K \pi_k \mathcal{N}(\mathbf{x} | \boldsymbol{\mu}_k, \boldsymbol{\Sigma}_k)$$

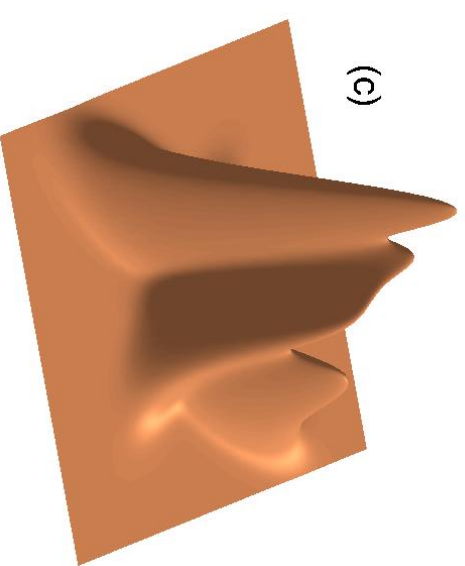
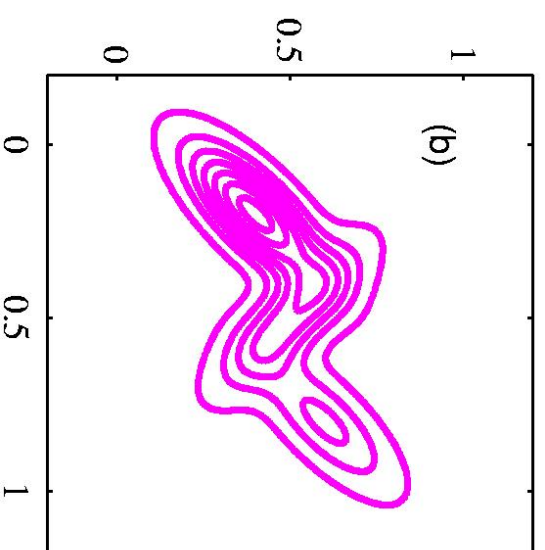
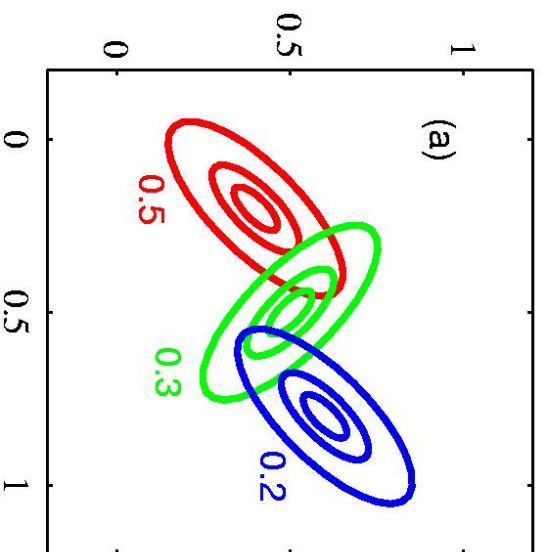
⏟
⏟
⏟

Mixing coefficient
Component

$$\forall k : \pi_k \geq 0 \quad \sum_{k=1}^K \pi_k = 1$$



Mixtures of Gaussians (Bishop 2.3.9)



Mixtures of Gaussians (Bishop 2.3.9)

- Determining parameters π , μ , and Σ using maximum log likelihood

$$\ln p(\mathbf{X} | \pi, \mu, \Sigma) = \sum_{n=1}^N \ln \left\{ \underbrace{\sum_{k=1}^K \pi_k \mathcal{N}(\mathbf{x}_n | \mu_k, \Sigma_k)}_{\text{Log of a sum; no closed form maximum.}} \right\}$$

- Solution: use standard, iterative, numeric optimization methods or the *expectation maximization* algorithm (Chapter 9).

Homework

Parametric Distributions

Basic building blocks:

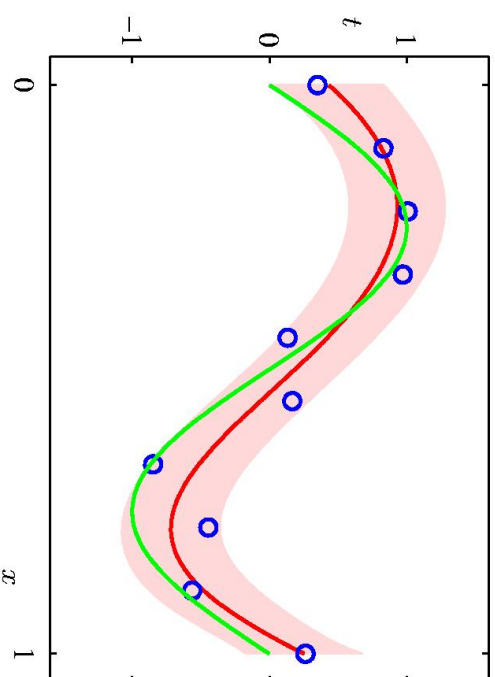
$$p(\mathbf{x}|\boldsymbol{\theta})$$

Need to determine $\boldsymbol{\theta}_{\text{given}}$ $\{\mathbf{x}_1, \dots, \mathbf{x}_N\}$

Representation: $\boldsymbol{\theta}^*$ or $p(\boldsymbol{\theta})$

Recall Curve Fitting

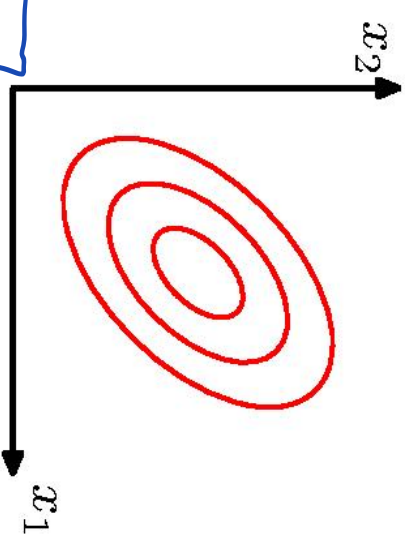
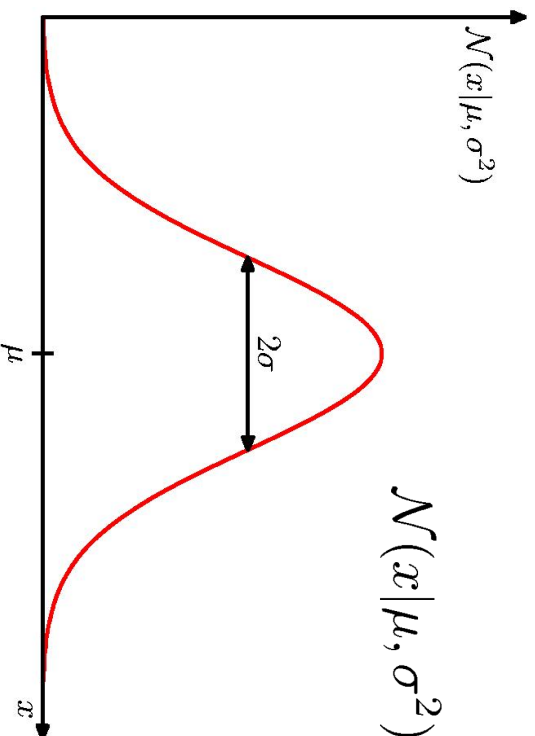
$$p(t|x, \mathbf{x}, \mathbf{t}) = \int p(t|x, \mathbf{w})p(\mathbf{w}|\mathbf{x}, \mathbf{t})d\mathbf{w}$$



We focus on Gaussians!

The Gaussian Distribution

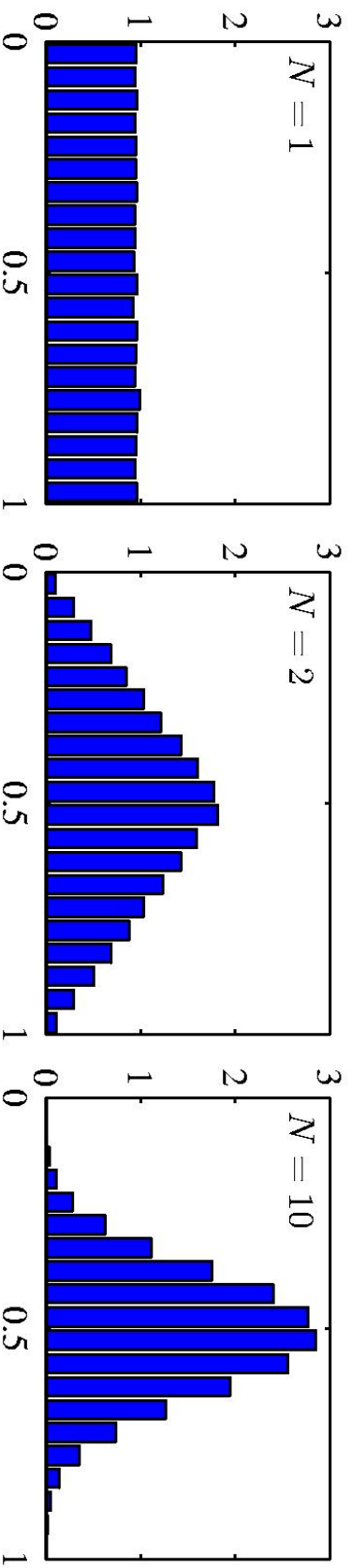
$$\mathcal{N}(x|\mu, \sigma^2) = \frac{1}{(2\pi\sigma^2)^{1/2}} \exp \left\{ -\frac{1}{2\sigma^2} (x - \mu)^2 \right\}$$



$$\mathcal{N}(\mathbf{x}|\boldsymbol{\mu}, \boldsymbol{\Sigma}) = \frac{1}{(2\pi)^{D/2}} \frac{1}{|\boldsymbol{\Sigma}|^{1/2}} \exp \left\{ -\frac{1}{2} (\mathbf{x} - \boldsymbol{\mu})^T \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu}) \right\}$$

Central Limit Theorem

- The distribution of the sum of N i.i.d. random variables becomes increasingly Gaussian as N grows.
- Example: N uniform $[0,1]$ random variables.



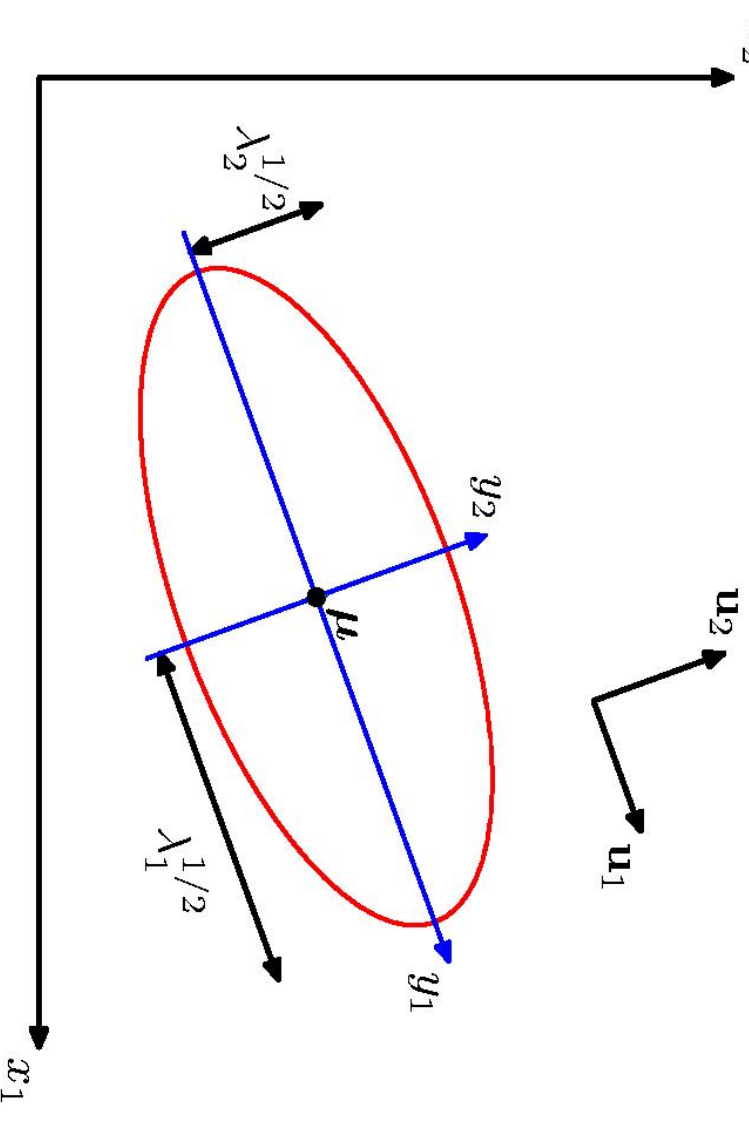
Geometry of the Multivariate Gaussian

$$\Delta^2 = (\mathbf{x} - \boldsymbol{\mu})^T \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu})$$

$$\boldsymbol{\Sigma}^{-1} = \sum_{i=1}^D \frac{1}{\lambda_i} \mathbf{u}_i \mathbf{u}_i^T$$

$$\Delta^2 = \sum_{i=1}^D \frac{y_i^2}{\lambda_i}$$

$$y_i = \mathbf{u}_i^T (\mathbf{x} - \boldsymbol{\mu})$$



Moments of the Multivariate Gaussian (2)

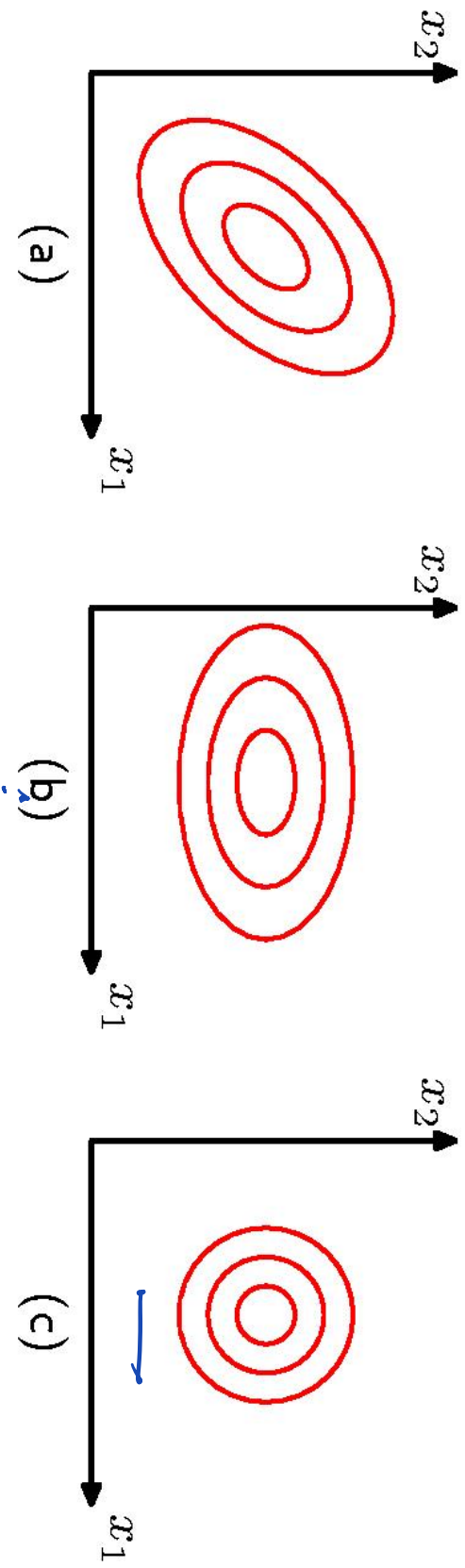
$$\mathbb{E}[\mathbf{x}\mathbf{x}^T] = \boldsymbol{\mu}\boldsymbol{\mu}^T + \boldsymbol{\Sigma}$$

$$\text{cov}[\mathbf{x}] = \mathbb{E}[(\mathbf{x} - \mathbb{E}[\mathbf{x}])(\mathbf{x} - \mathbb{E}[\mathbf{x}])^T] = \boldsymbol{\Sigma}$$

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

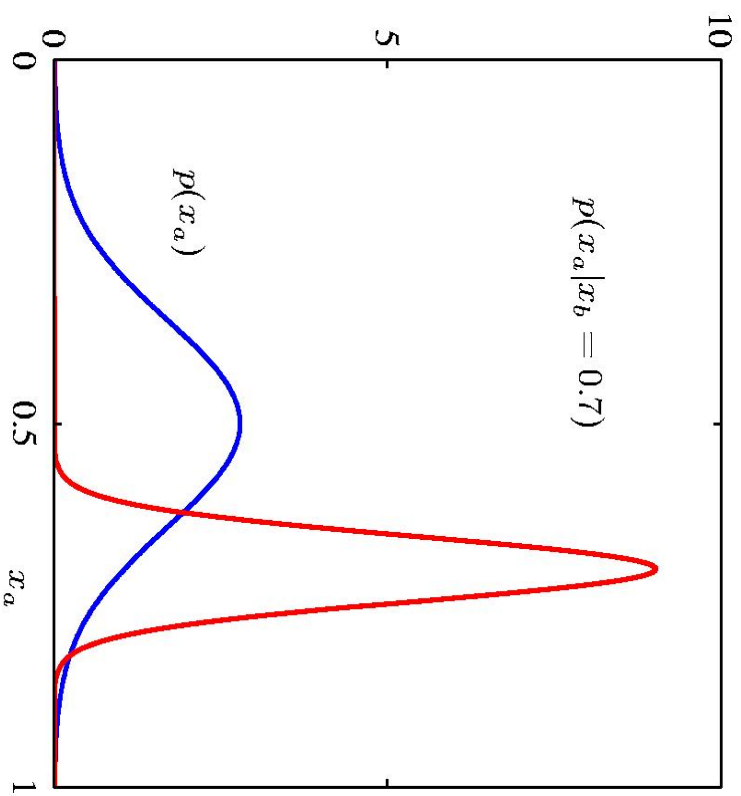
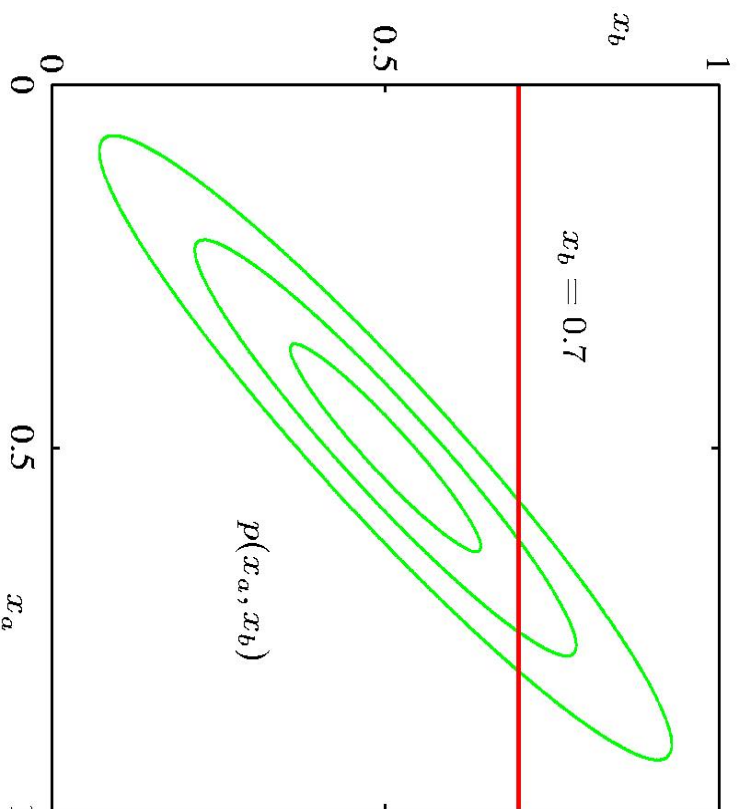
$$= \sigma^2 \mathbf{I}$$

A Gaussian requires $D^*(D-1)/2 + D$ parameters.
Often we use $D + D$ or
Just $D+1$ parameters.



Partitioned Conditionals and Marginals, page 89

$$\begin{aligned}
 p(\mathbf{x}_a | \mathbf{x}_b) &= \mathcal{N}(\mathbf{x}_a | \underline{\mu}_{a|b}, \underline{\Sigma}_{a|b}) \\
 \Sigma_{a|b} &= \Lambda_{aa}^{-1} = \Sigma_{aa} - \Sigma_{ab} \Sigma_{bb}^{-1} \Sigma_{ba} \\
 \mu_{a|b} &= \Sigma_{a|b} \{ \Lambda_{aa} \mu_a - \Lambda_{ab} (\mathbf{x}_b - \mu_b) \} \\
 &= \mu_a - \Lambda_{aa}^{-1} \Lambda_{ab} (\mathbf{x}_b - \mu_b) \\
 &= \mu_a + \Sigma_{ab} \Sigma_{bb}^{-1} (\mathbf{x}_b - \mu_b)
 \end{aligned}
 \qquad
 \begin{aligned}
 p(\mathbf{x}_a) &= \int p(\mathbf{x}_a, \mathbf{x}_b) d\mathbf{x}_b \\
 &= \mathcal{N}(\mathbf{x}_a | \underline{\mu}_a, \underline{\Sigma}_{aa})
 \end{aligned}$$



Maximum Likelihood for the Gaussian (1)

Given i.i.d. data $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_N)^T$, the log likelihood function is given by

$$\ln p(\mathbf{X} | \mu, \Sigma) = -\frac{ND}{2} \ln(2\pi) - \frac{N}{2} \ln |\Sigma| - \frac{1}{2} \sum_{n=1}^N (\mathbf{x}_n - \mu)^T \Sigma^{-1} (\mathbf{x}_n - \mu)$$

$$\ell = -\log p(\mathbf{X} | \mu, \Sigma) = \frac{N}{2} \log |\Sigma| + \frac{1}{2} \text{tr} \left(\sum_n (\mathbf{x}_n - \mu)(\mathbf{x}_n - \mu)^T \Sigma^{-1} \right)$$

$$\frac{\partial \ell}{\partial \mu} = -\sum_n \Sigma^{-1} (\mathbf{x}_n - \mu) = 0$$

$$\mu_{ML} = \frac{\Sigma}{N} \sum_n \mathbf{x}_n$$

$$\Sigma_y = \frac{1}{N} \sum_n (\mathbf{x}_n - \mu)(\mathbf{x}_n - \mu)^T$$

$$\frac{\partial \ell}{\partial \Sigma} = \Sigma^{-1} - \sum_y \Sigma^{-1} \Sigma^{-1} = 0 \quad \Sigma_{ML} = \Sigma_y$$

$$\frac{\partial}{\partial \mathbf{A}} \ln |\mathbf{A}| = (\mathbf{A}^{-1})^T \quad (\text{C.28})$$

$$\frac{\partial}{\partial \mathbf{A}} \text{Tr}(\mathbf{A}\mathbf{B}) = \mathbf{B}^T \quad (\text{C.24})$$

$$\frac{\partial}{\partial x} (\mathbf{A}^{-1}) = -\mathbf{A}^{-1} \frac{\partial \mathbf{A}}{\partial x} \mathbf{A}^{-1} \quad (\text{C.21})$$

Maximum Likelihood for the Gaussian

- Set the derivative of the log likelihood function to zero,

$$\frac{\partial}{\partial \boldsymbol{\mu}} \ln p(\mathbf{X} | \boldsymbol{\mu}, \boldsymbol{\Sigma}) = \sum_{n=1}^N \boldsymbol{\Sigma}^{-1} (\mathbf{x}_n - \boldsymbol{\mu}) = 0$$

- and solve to obtain

$$\boldsymbol{\mu}_{\text{ML}} = \frac{1}{N} \sum_{n=1}^N \mathbf{x}_n.$$

- Similarly

$$\boldsymbol{\Sigma}_{\text{ML}} = \frac{1}{N} \sum_{n=1}^N (\mathbf{x}_n - \boldsymbol{\mu}_{\text{ML}})(\mathbf{x}_n - \boldsymbol{\mu}_{\text{ML}})^{\text{T}}.$$