

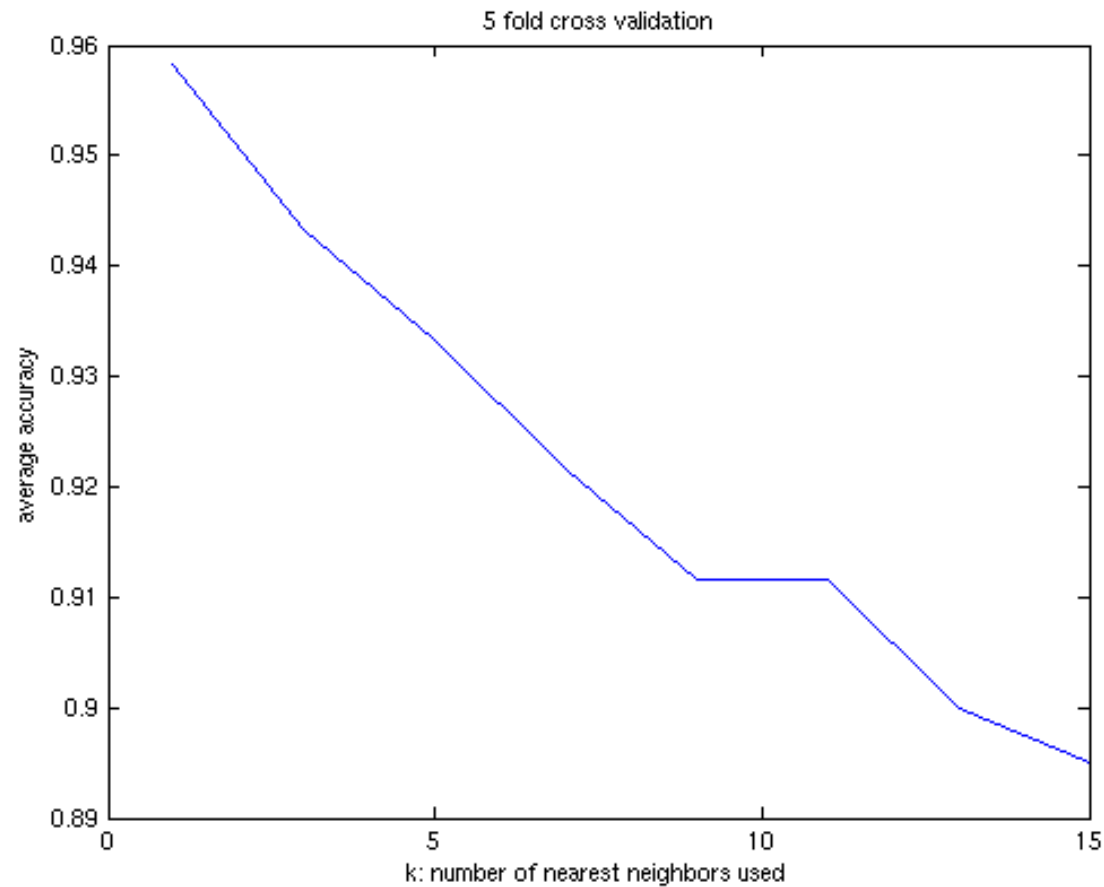
Introduction to Machine Learning

Brown University CSCI 1950-F, Spring 2011
Prof. Erik Sudderth

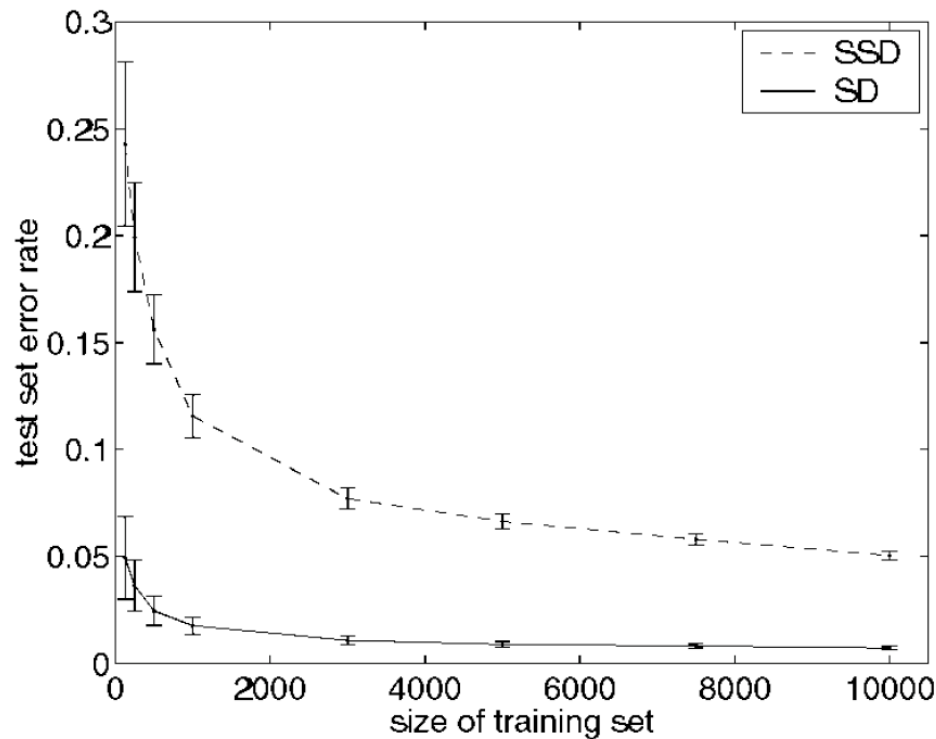
Lecture 9: Bayesian Regression,
Logistic Regression

Many figures courtesy Kevin Murphy's textbook,
Machine Learning: A Probabilistic Perspective

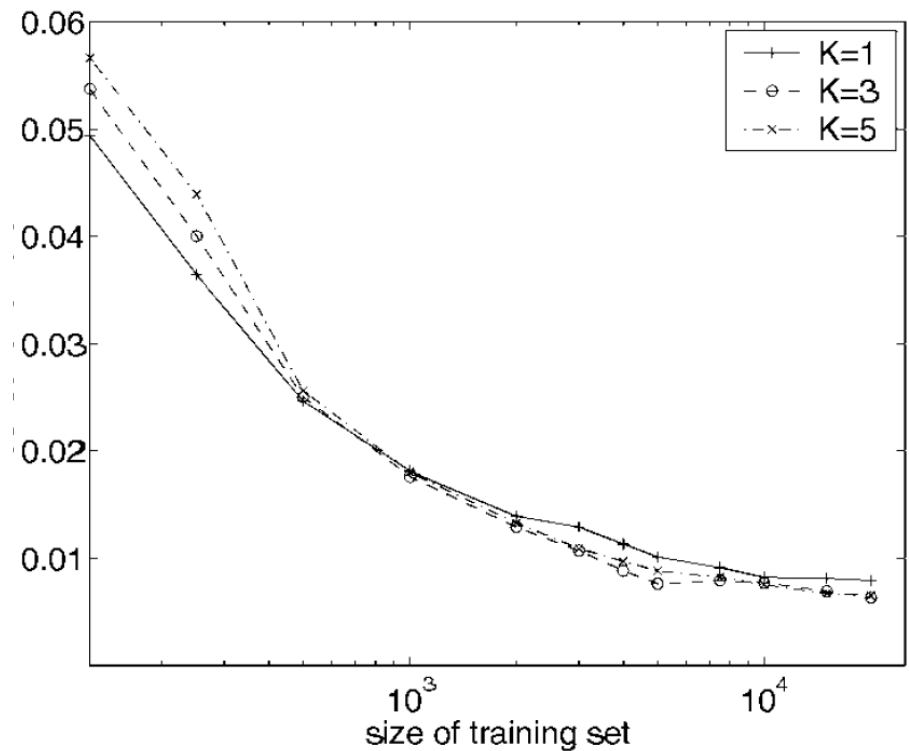
K-NN Cross-Validation: MNIST Digits (Homework 3)



K-NN Performance: Shape Contexts

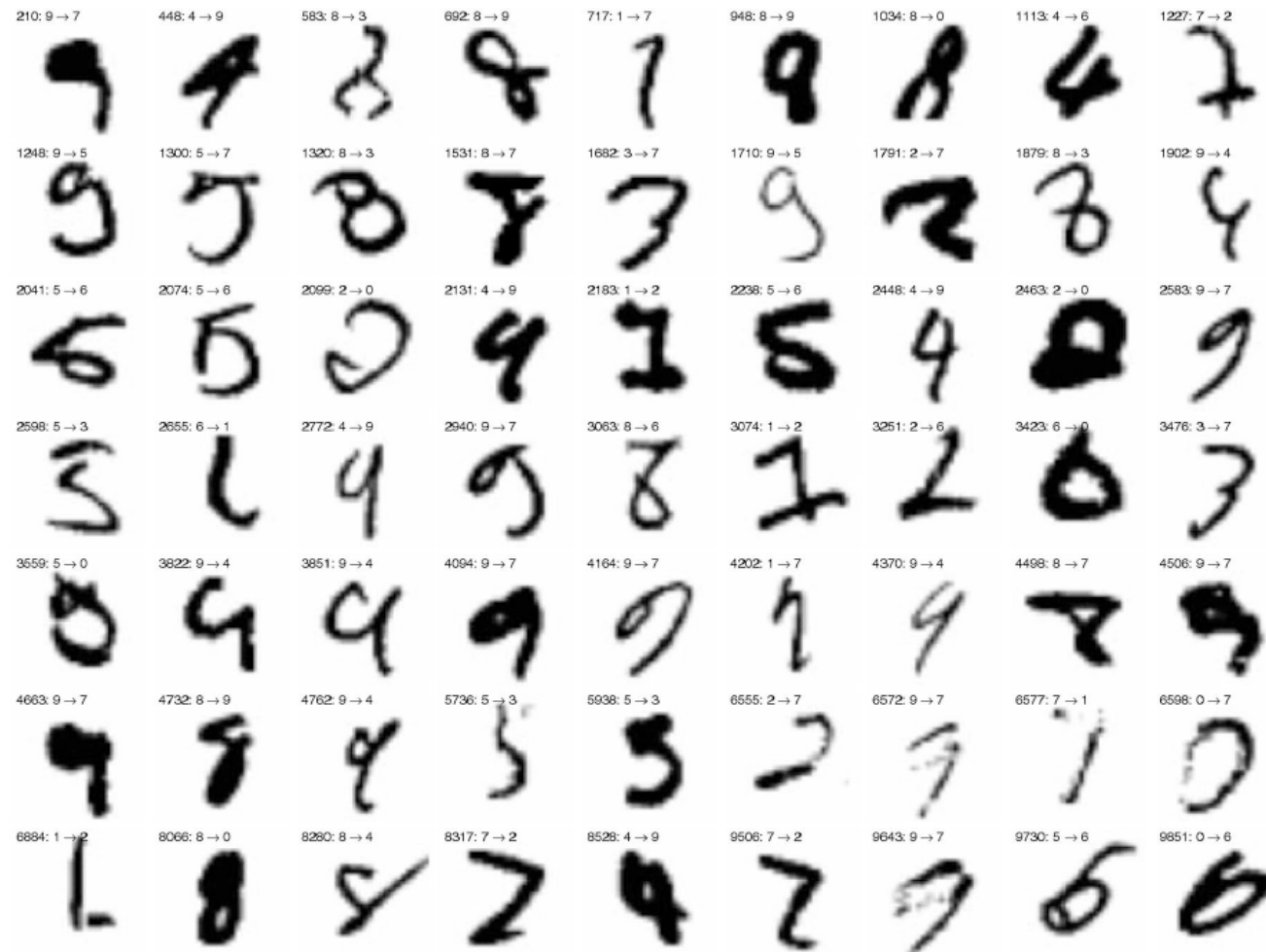


Alternative Distance Measures

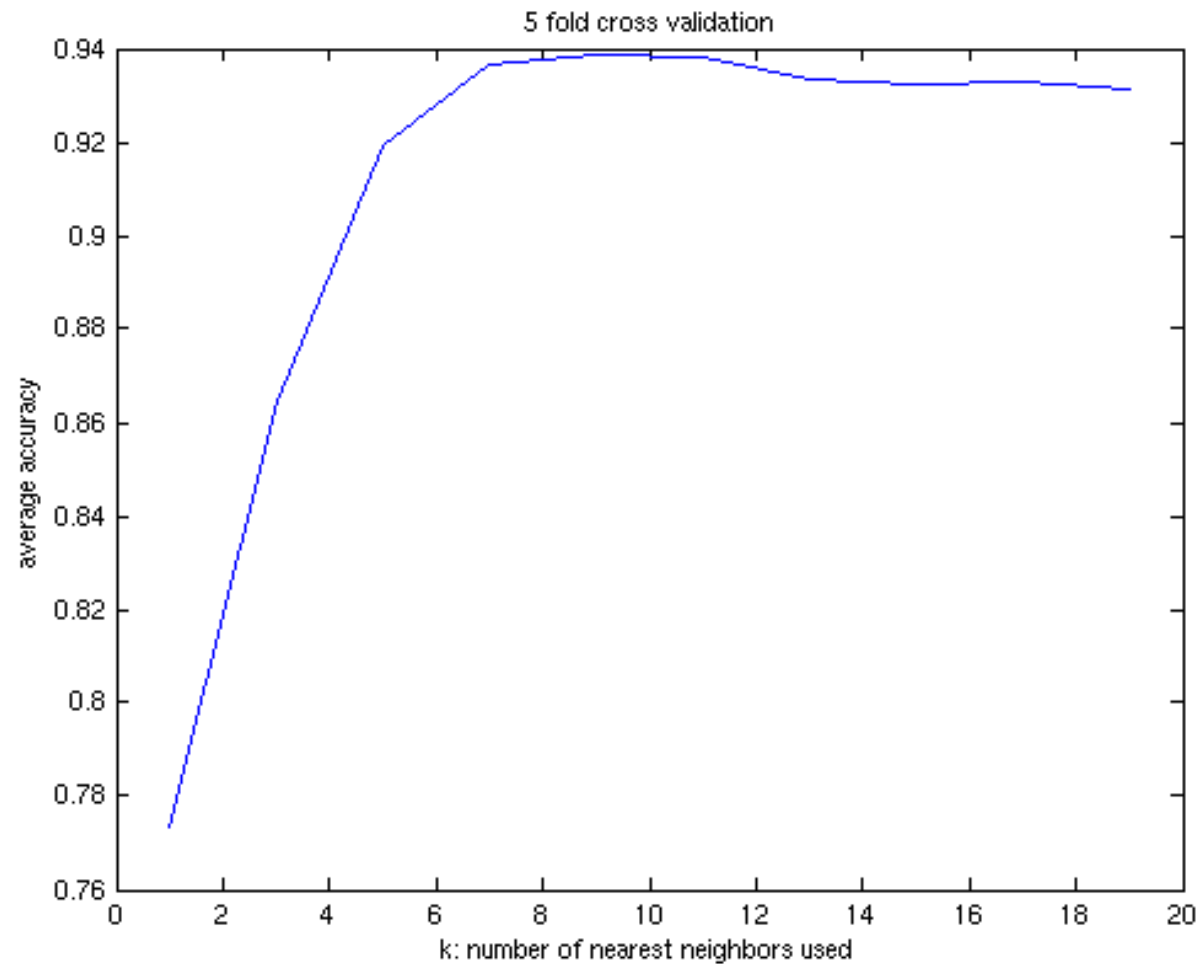


Choice of Neighborhood Size K

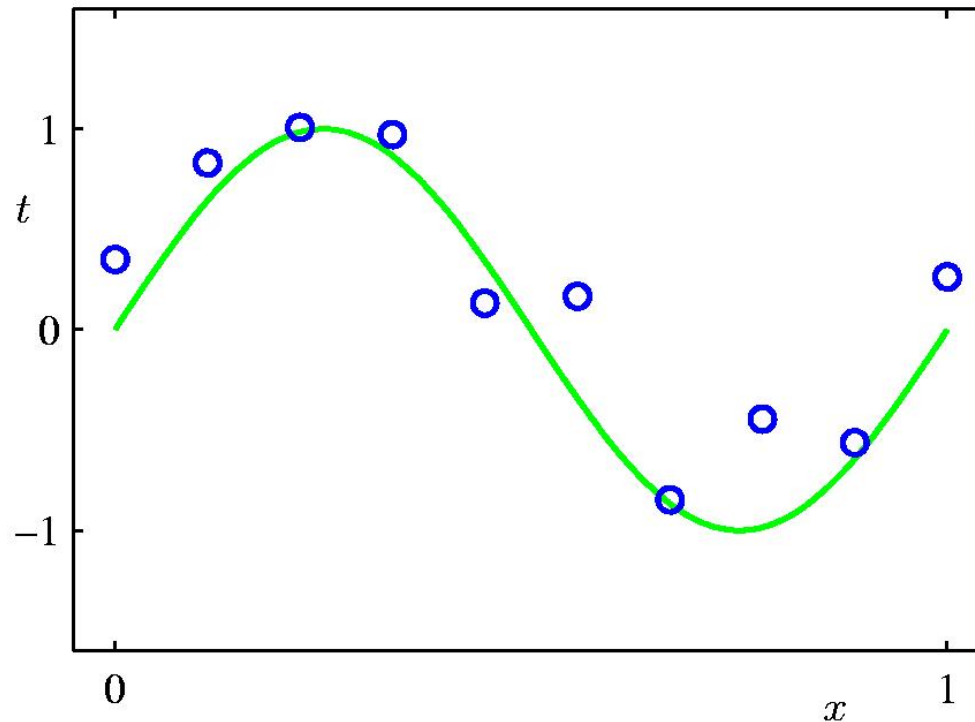
MNIST: Shape Context Errors



K-NN Cross-Validation: Nursery



Polynomial Curve Fitting



$$y(x, \mathbf{w}) = w_0 + w_1x + w_2x^2 + \dots + w_Mx^M = \sum_{j=0}^M w_jx^j$$

*Slides adapted from Bishop's "Pattern Recognition and Machine Learning"
Notation differs from Murphy's "Machine Learning: A Probabilistic Perspective"*

Regularized Least Squares (1)

- Consider the error function:

$$E_D(\mathbf{w}) + \lambda E_W(\mathbf{w})$$

Data term + Regularization term

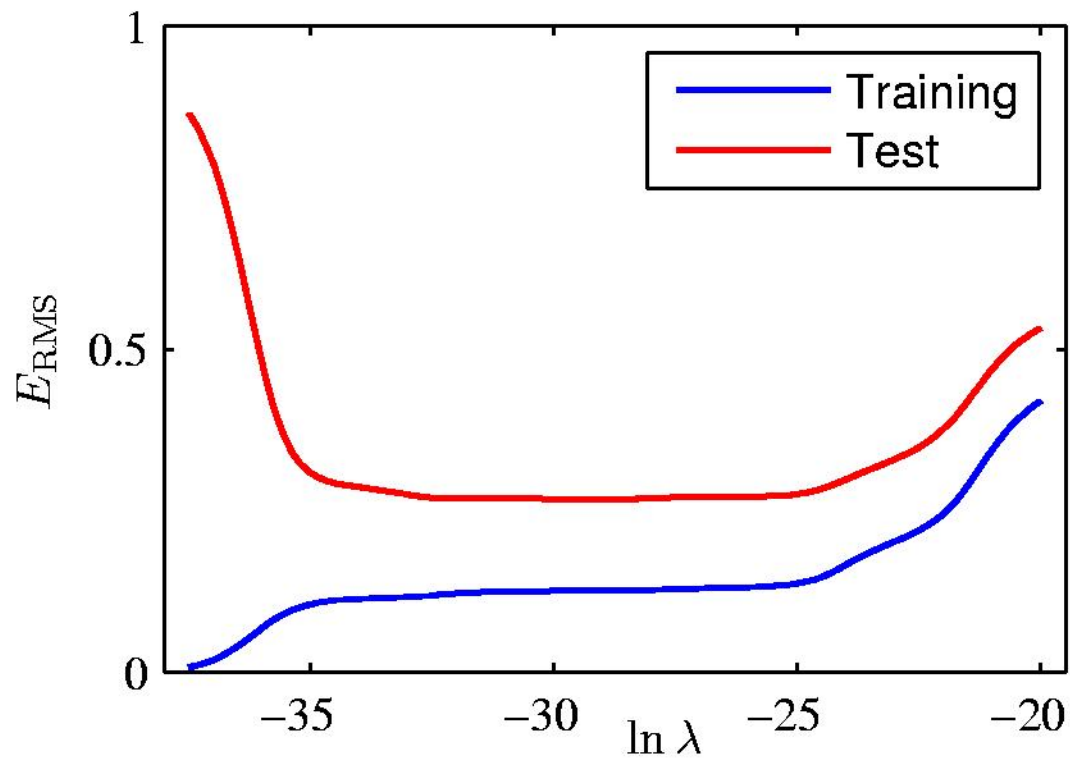
- With the sum-of-squares error function and a quadratic regularizer, we get

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

- which is minimized by

$$\mathbf{w} = \left(\lambda \mathbf{I} + \Phi^T \Phi \right)^{-1} \Phi^T \mathbf{t}.$$

Regularization: E_{RMS} vs. $\ln \lambda$



Bayesian Linear Regression

- Gaussian prior and likelihood imply Gaussian posteriors:

$$p(\mathbf{w}|\mathbf{t}) = \mathcal{N}(\mathbf{w}|\mathbf{m}_N, \mathbf{S}_N)$$

$$p(\mathbf{w}) = \mathcal{N}(\mathbf{w}|\mathbf{m}_0, \mathbf{S}_0).$$

$$\mathbf{m}_N = \mathbf{S}_N \left(\mathbf{S}_0^{-1} \mathbf{m}_0 + \beta \Phi^T \mathbf{t} \right)$$

$$\mathbf{S}_N^{-1} = \mathbf{S}_0^{-1} + \beta \Phi^T \Phi.$$

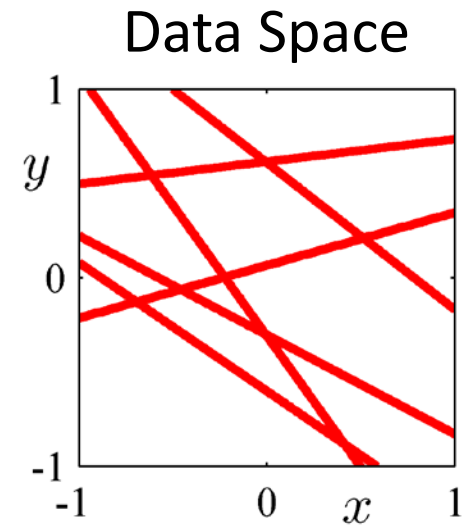
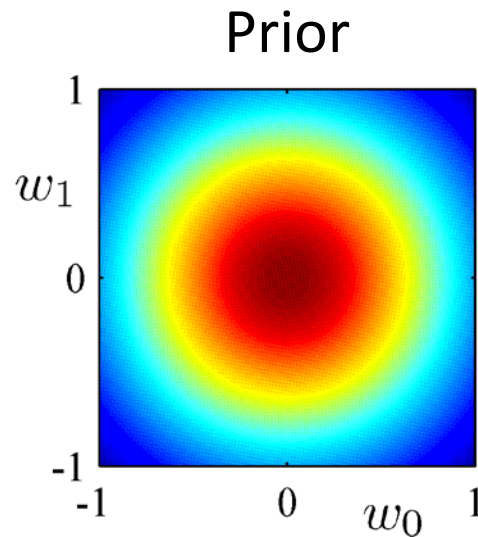
$$p(\mathbf{w}) = \mathcal{N}(\mathbf{w}|\mathbf{0}, \alpha^{-1} \mathbf{I})$$

$$\mathbf{m}_N = \beta \mathbf{S}_N \Phi^T \mathbf{t}$$

$$\mathbf{S}_N^{-1} = \alpha \mathbf{I} + \beta \Phi^T \Phi.$$

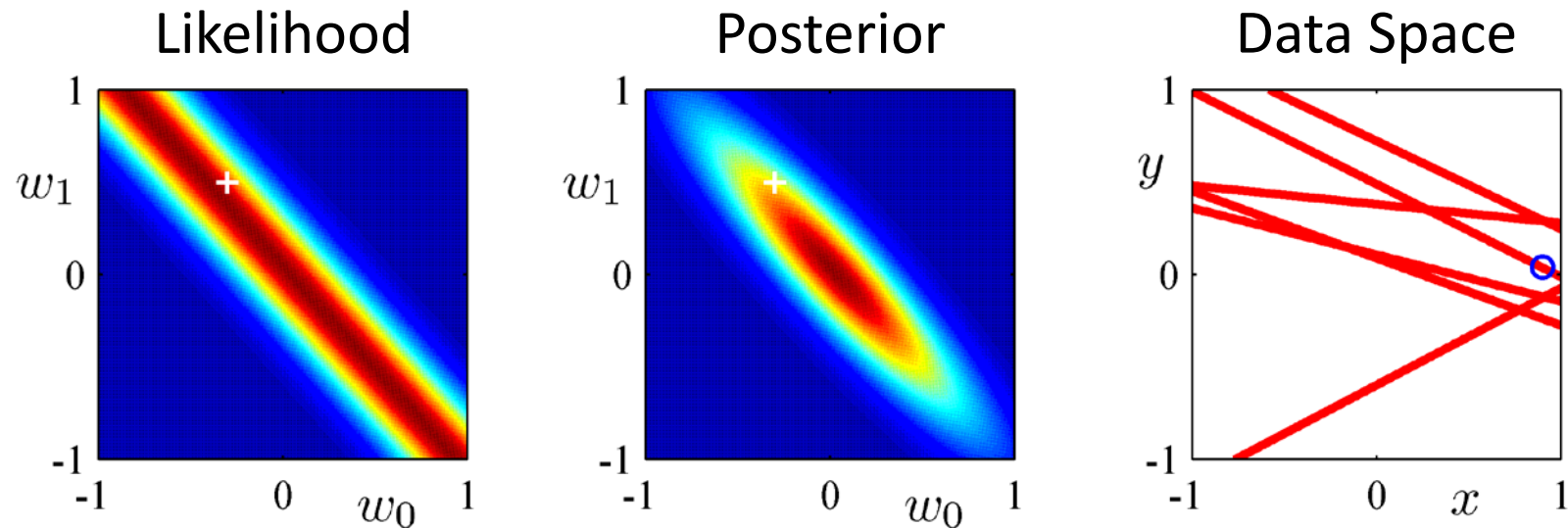
Bayesian Linear Regression (3)

0 data points observed



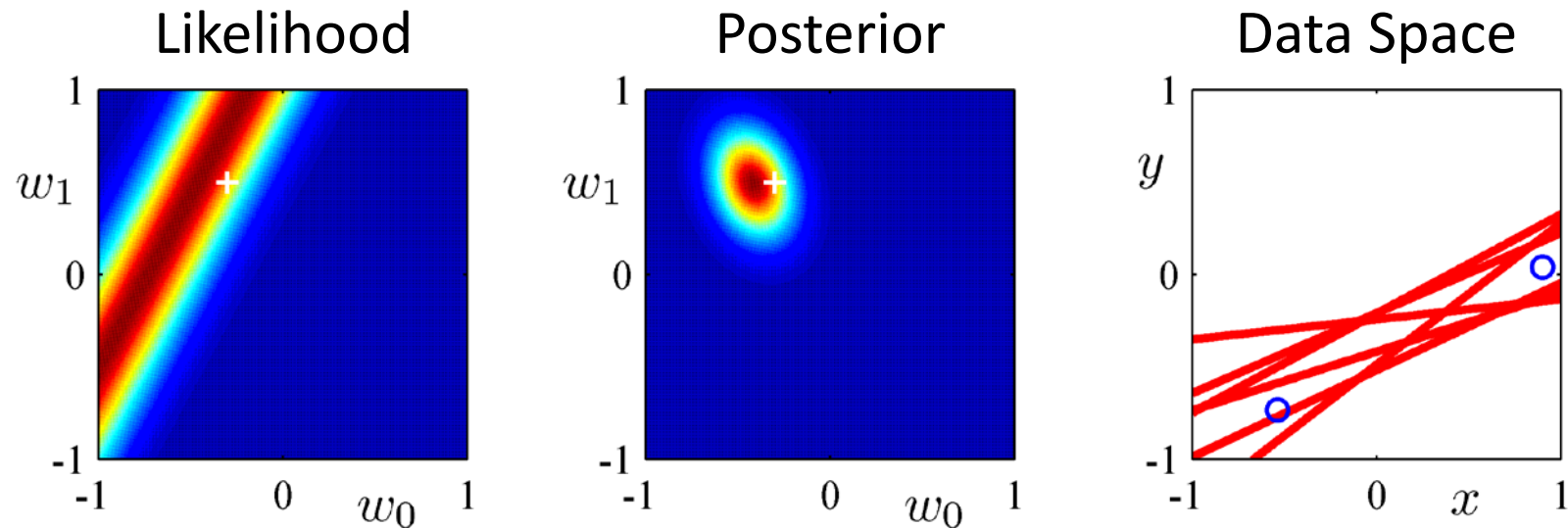
Bayesian Linear Regression (4)

1 data point observed



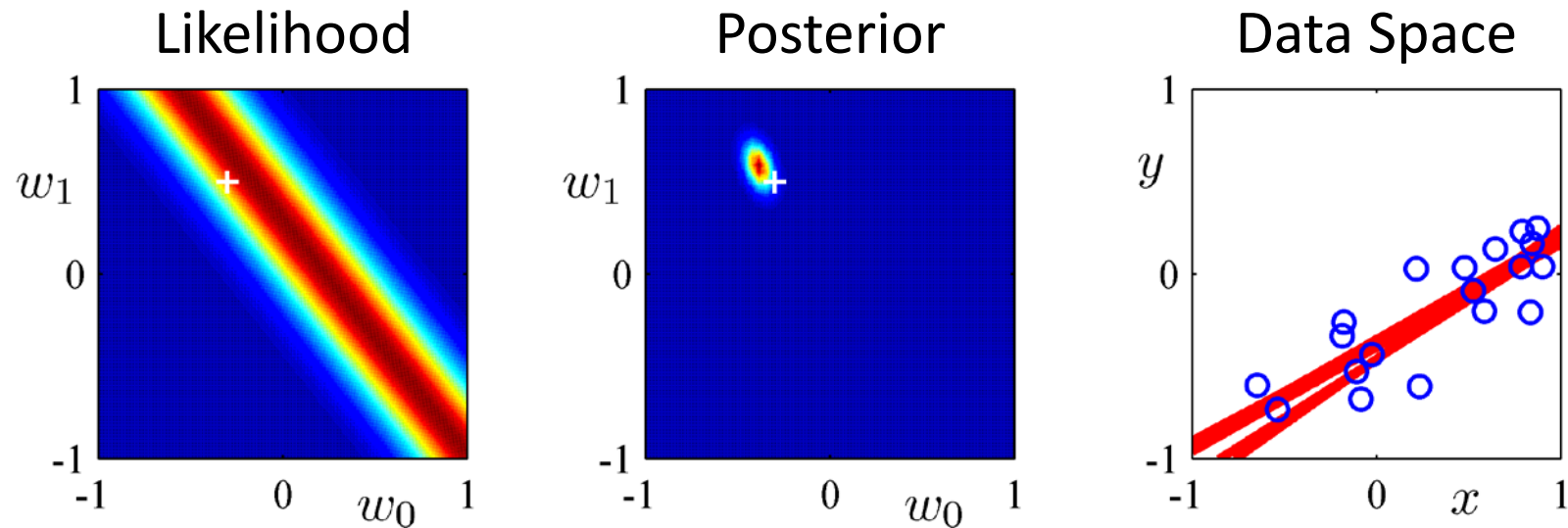
Bayesian Linear Regression (5)

2 data points observed



Bayesian Linear Regression (6)

20 data points observed



Predictive Distribution (1)

- Predict t for new values of $\mathbf{x}$ by integrating over $\mathbf{w}$:

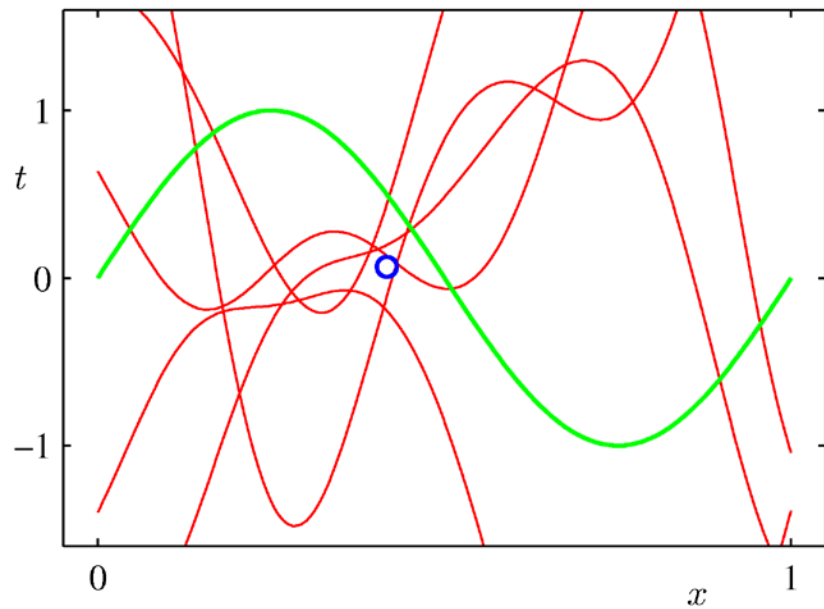
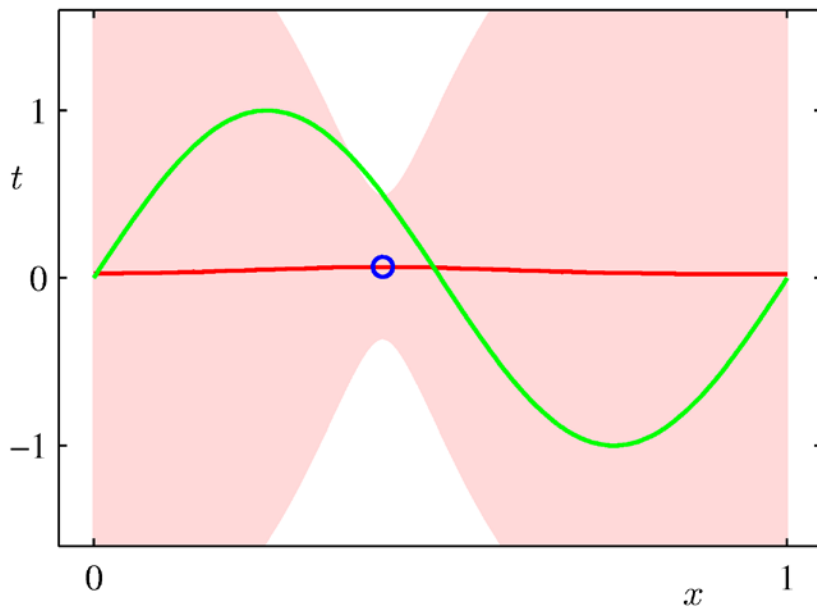
$$\begin{aligned} p(t|\mathbf{t}, \alpha, \beta) &= \int p(t|\mathbf{w}, \beta) p(\mathbf{w}|\mathbf{t}, \alpha, \beta) d\mathbf{w} \\ &= \mathcal{N}(t|\mathbf{m}_N^T \phi(\mathbf{x}), \sigma_N^2(\mathbf{x})) \end{aligned}$$

- where

$$\sigma_N^2(\mathbf{x}) = \frac{1}{\beta} + \phi(\mathbf{x})^T \mathbf{S}_N \phi(\mathbf{x}).$$

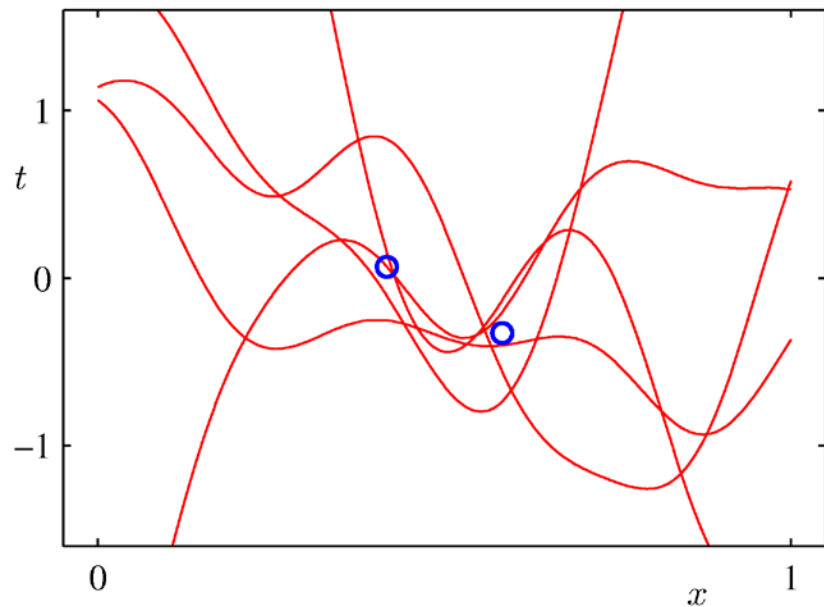
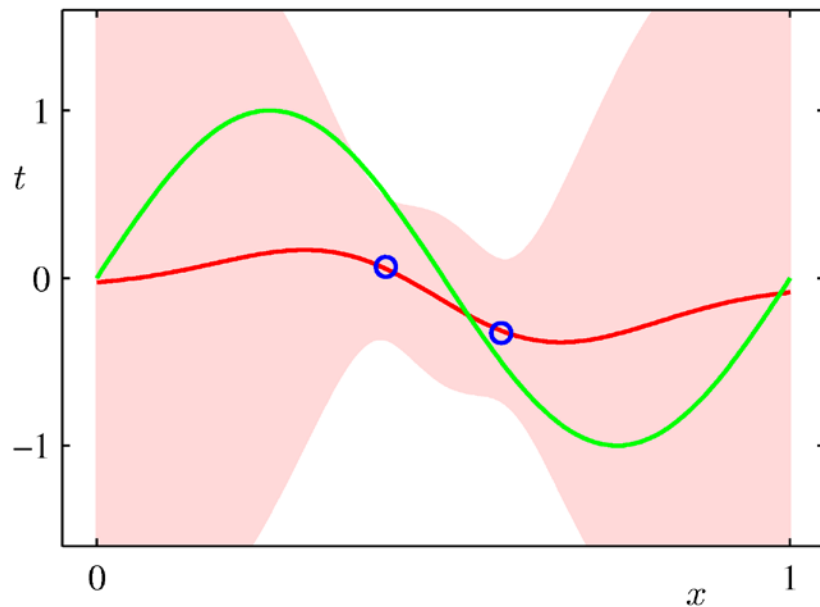
Predictive Distribution (2)

- Example: Sinusoidal data, 9 Gaussian basis functions, 1 data point



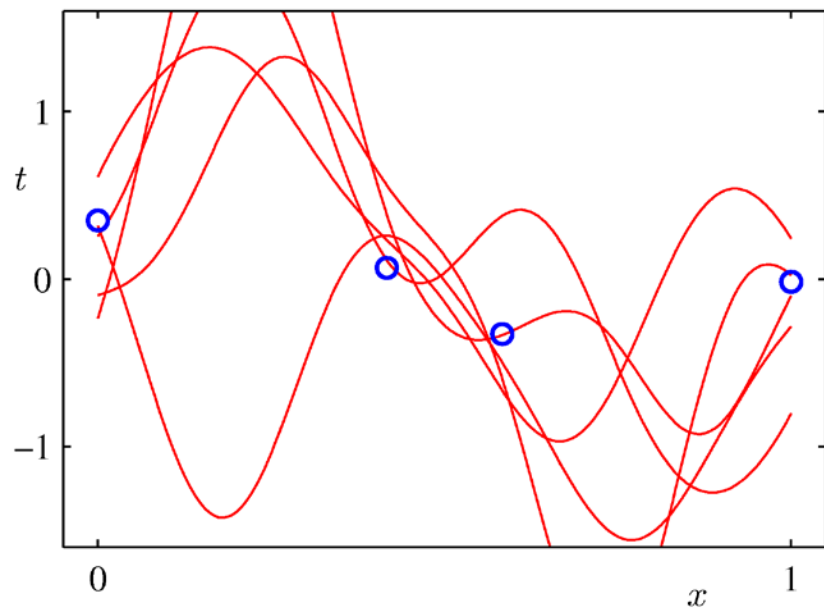
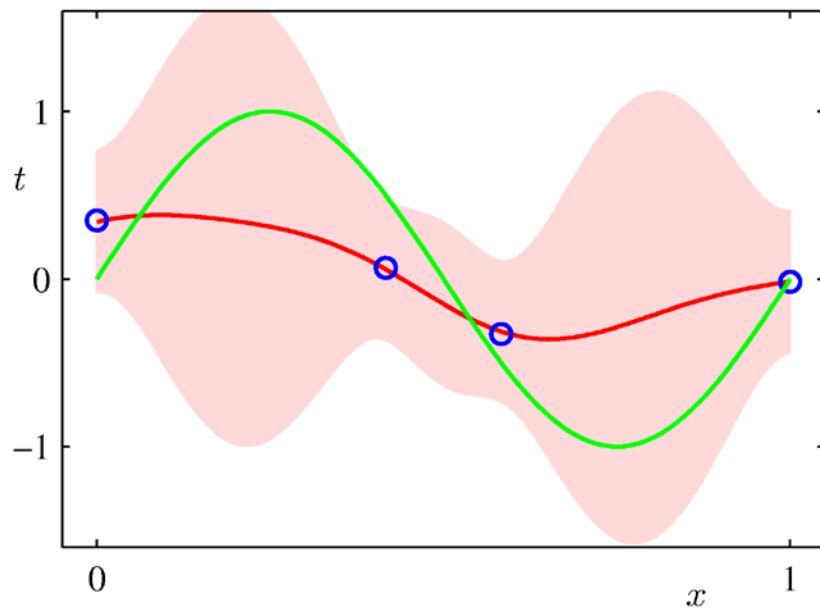
Predictive Distribution (3)

- Example: Sinusoidal data, 9 Gaussian basis functions, 2 data points



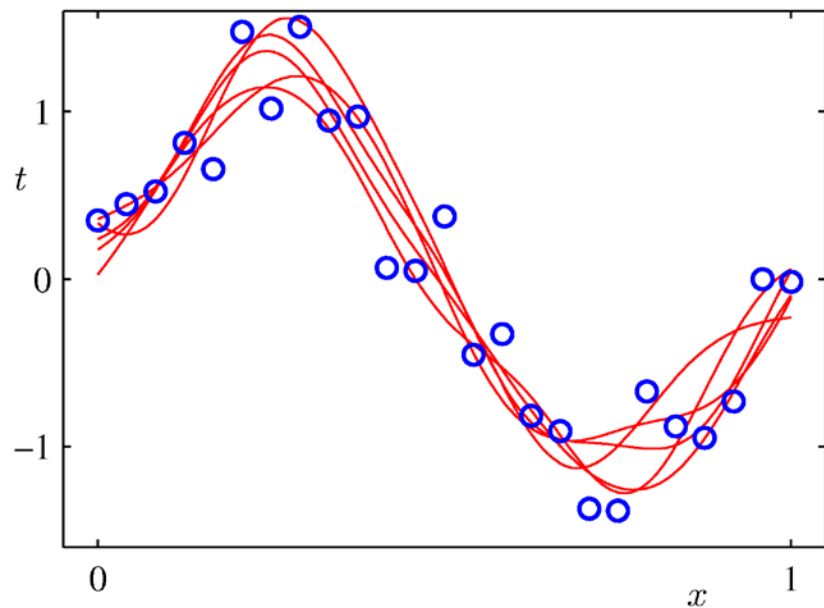
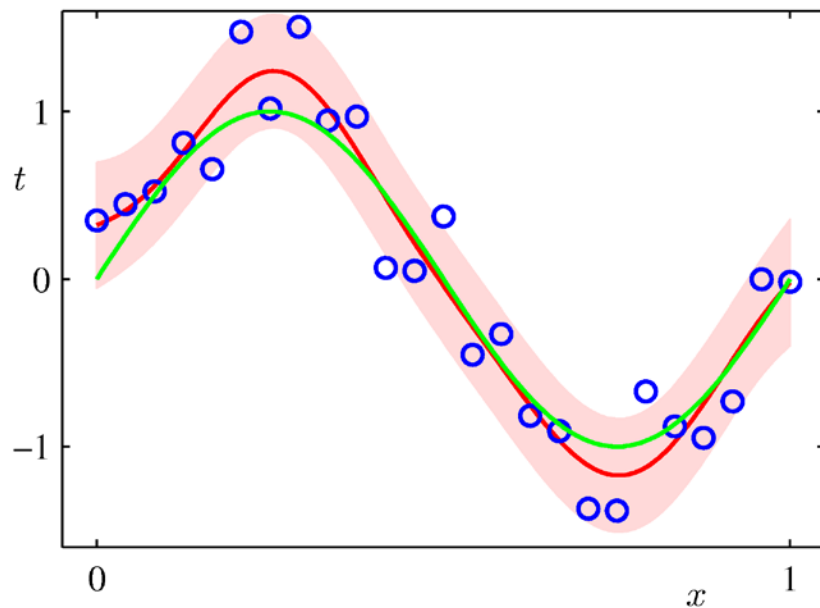
Predictive Distribution (4)

- Example: Sinusoidal data, 9 Gaussian basis functions, 4 data points

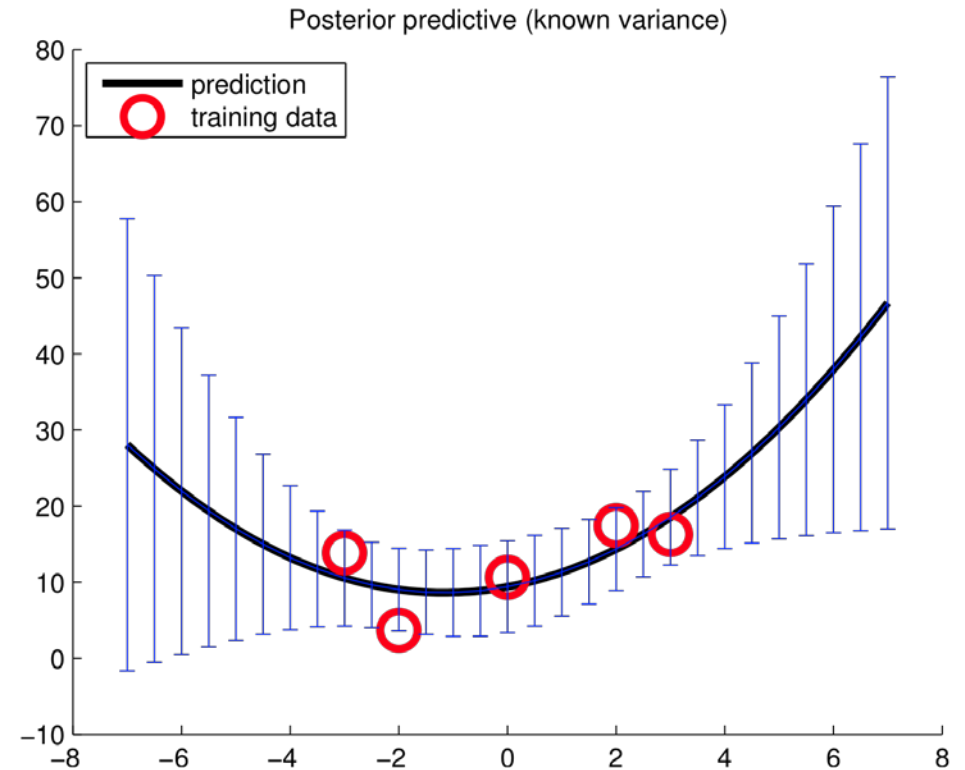
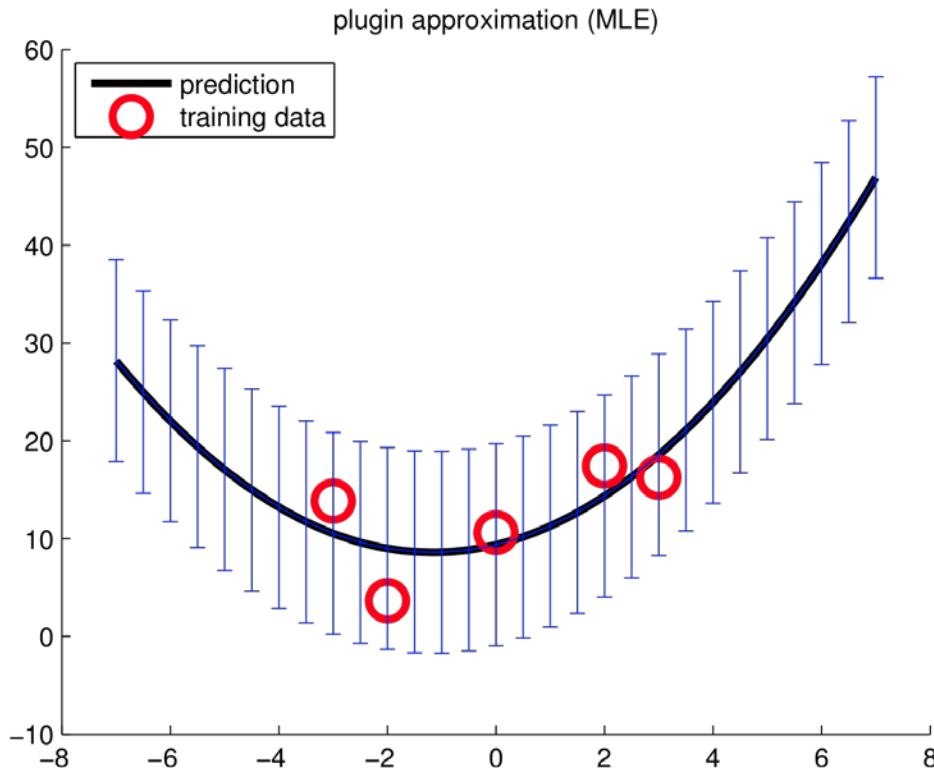


Predictive Distribution (5)

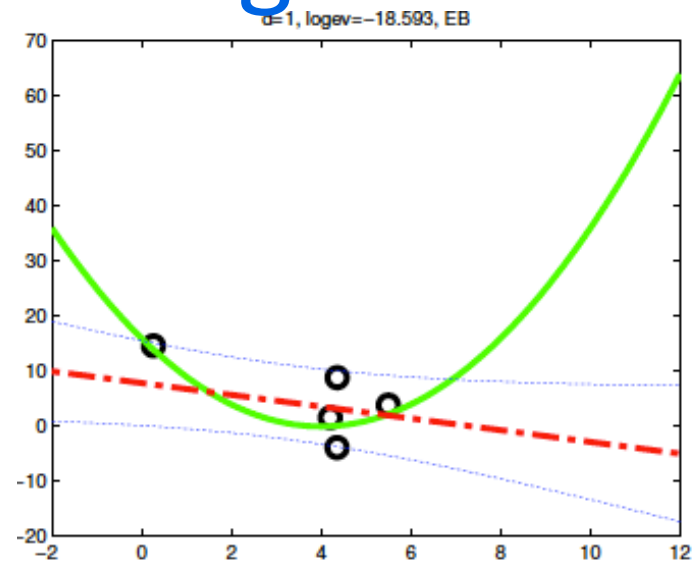
- Example: Sinusoidal data, 9 Gaussian basis functions, 25 data points



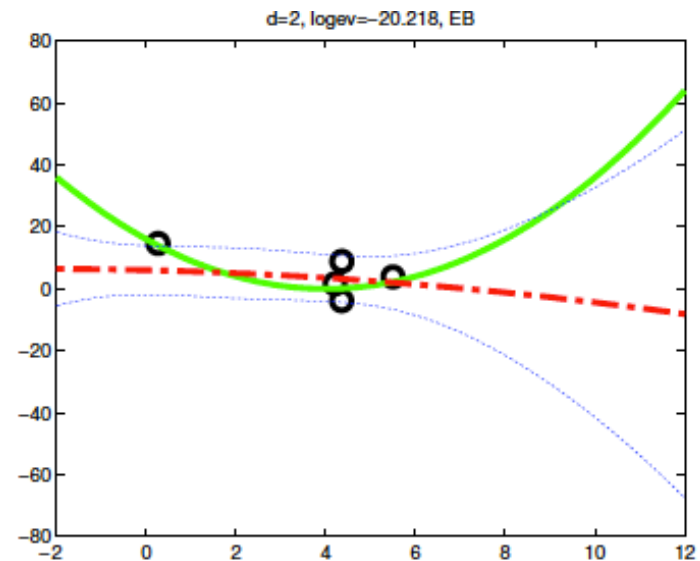
Estimation vs. Predictive Distributions



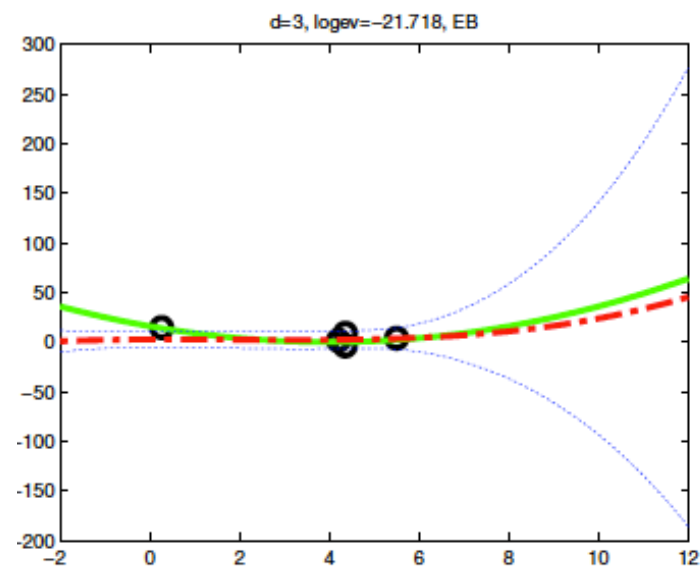
Marginal Data Likelihood: N=5



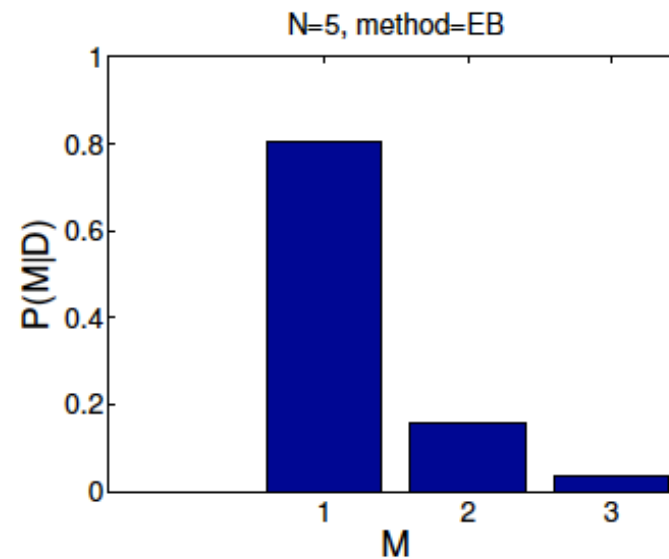
(a)



(b)

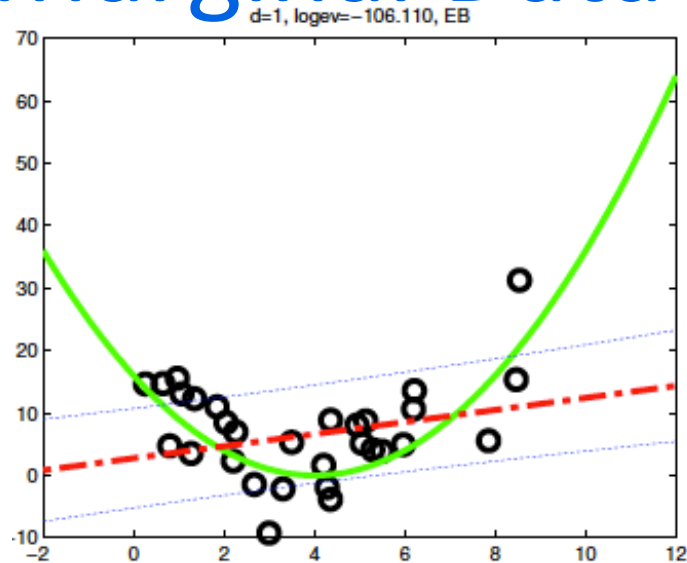


(c)

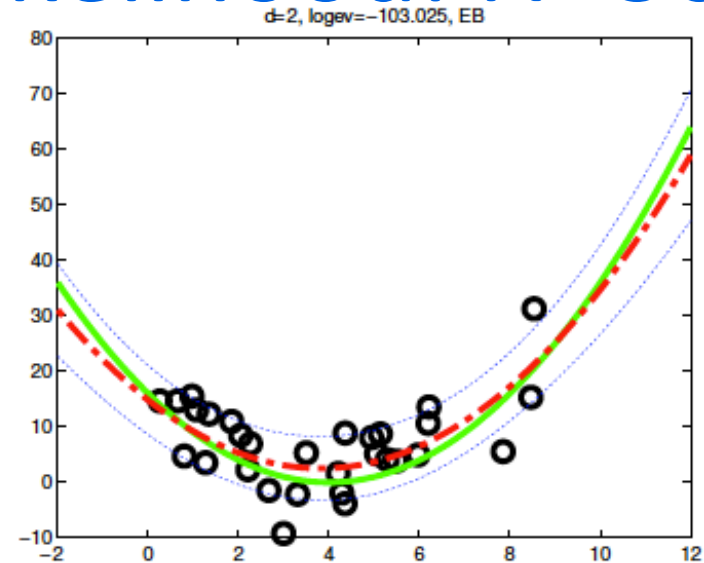


(d)

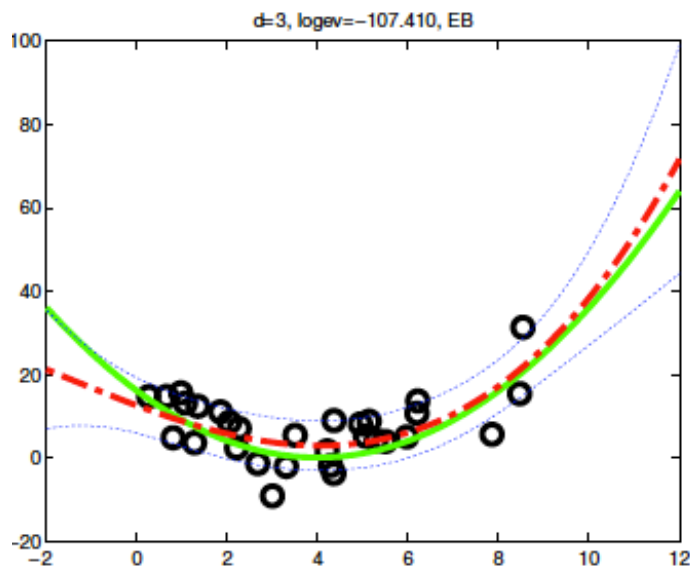
Marginal Data Likelihood: $N=30$



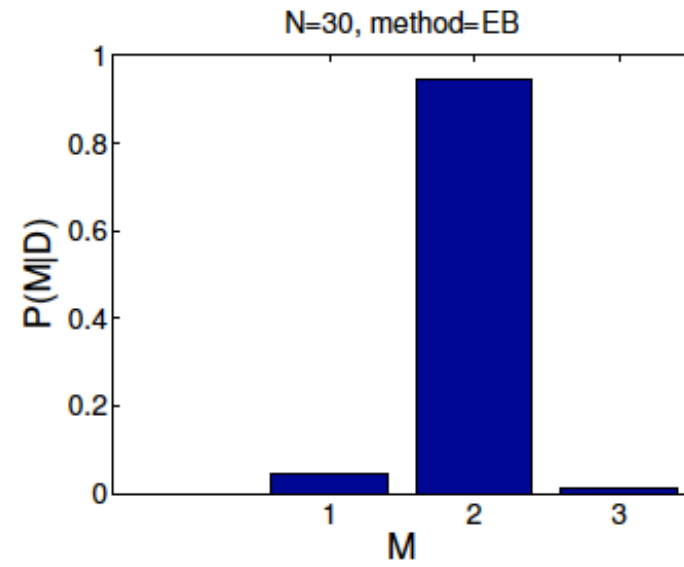
(a)



(b)

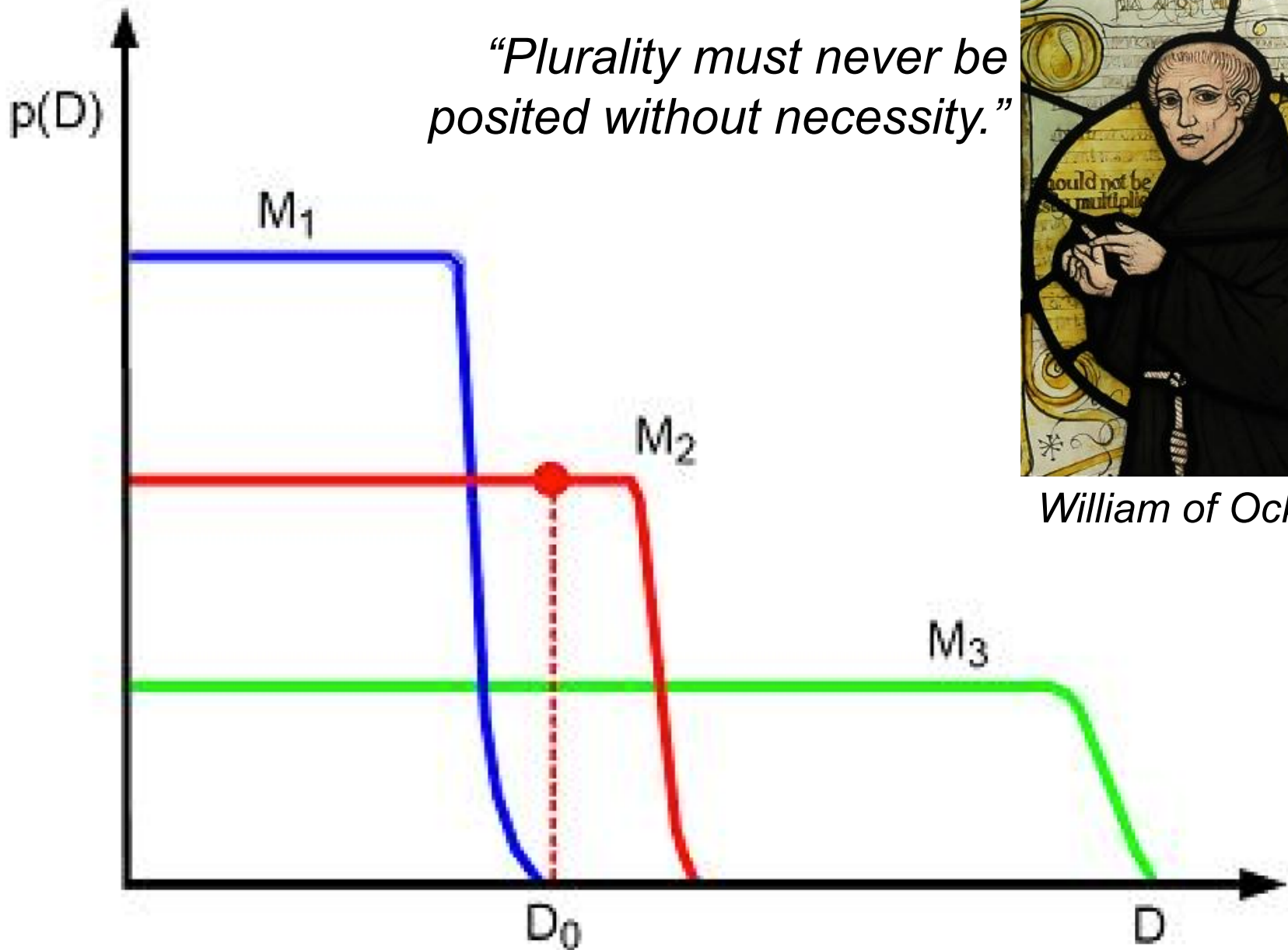


(c)



(d)

Bayesian Ockham's Razor



William of Ockham

Back to Classification

Supervised Learning

Unsupervised Learning

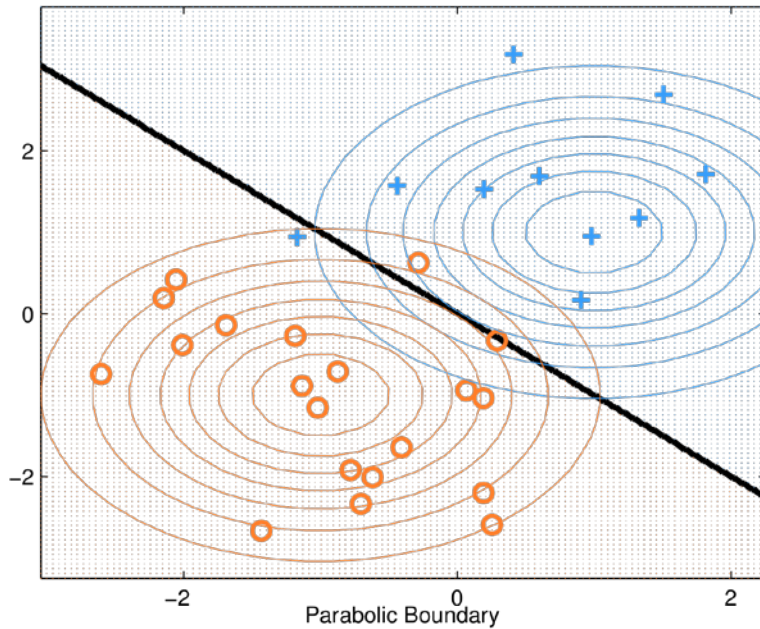
Discrete
Continuous

classification or categorization	clustering
regression	dimensionality reduction

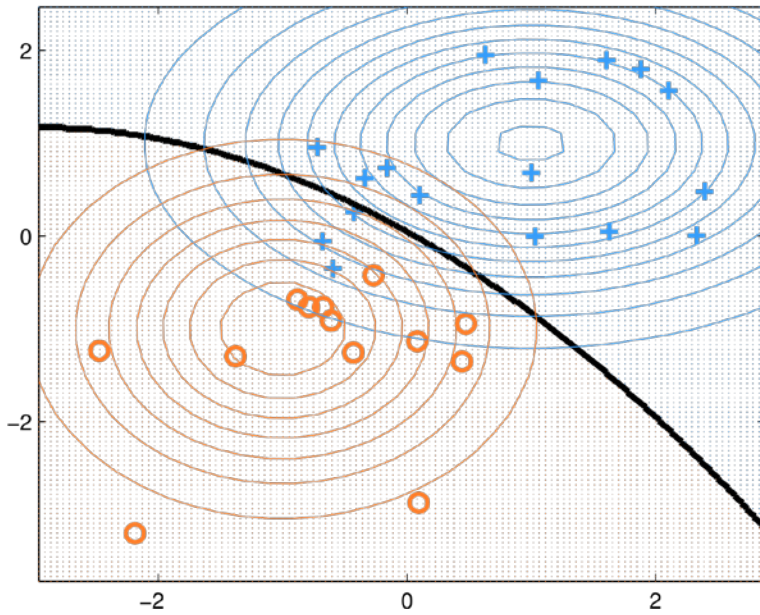
- Predict label/response y from feature x
- Generative classification: Apply Bayes' rule to learned $p(x,y)$
- Discriminative or conditional regression & classification: directly learn a model of $p(y | x)$, assuming x always given

Discriminant Analysis

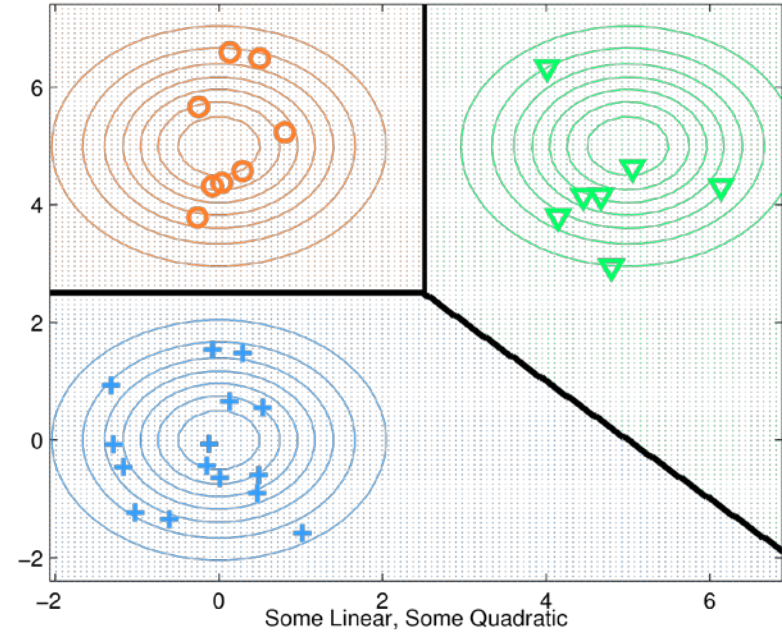
Linear Boundary



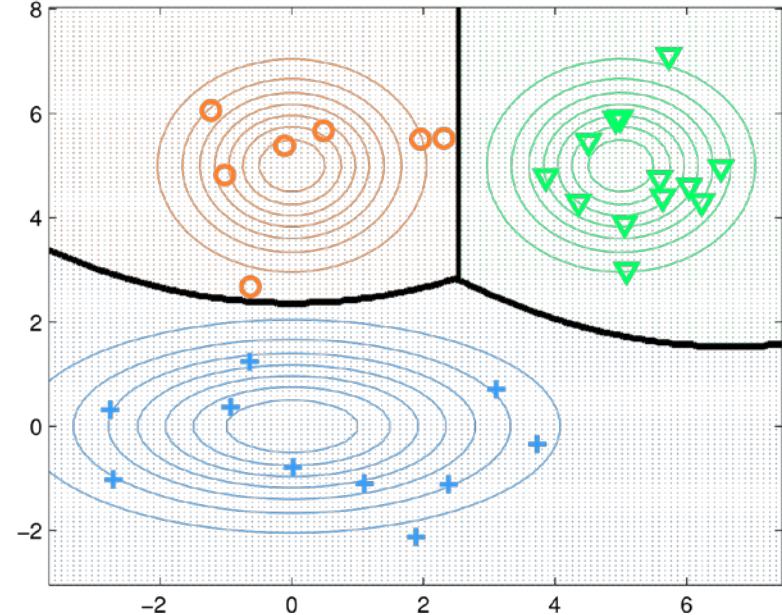
Parabolic Boundary



All Linear Boundaries

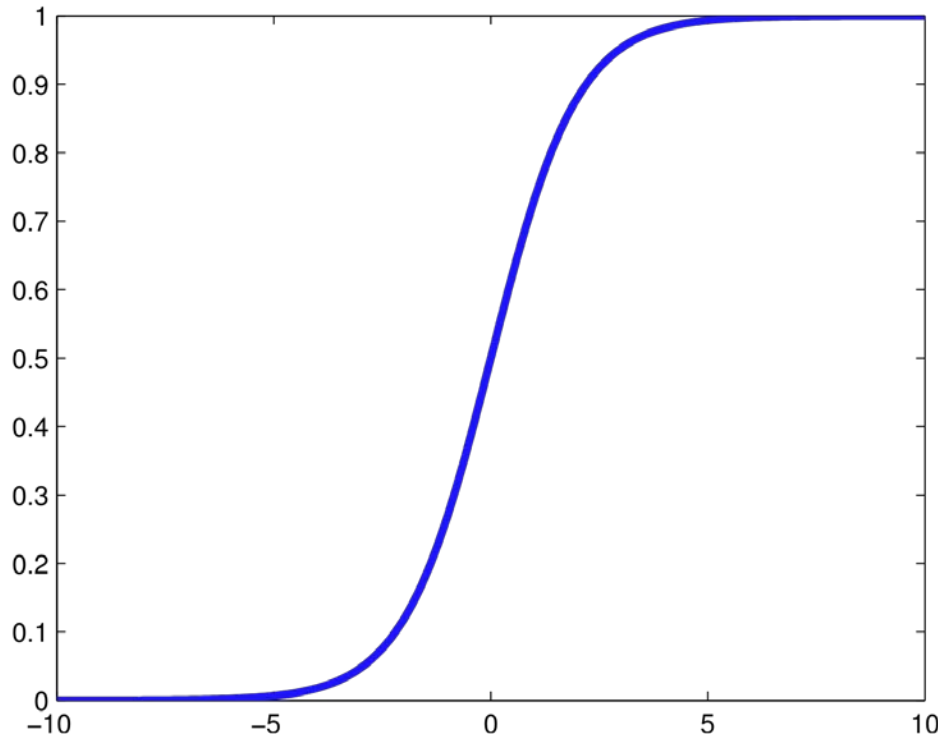


Some Linear, Some Quadratic

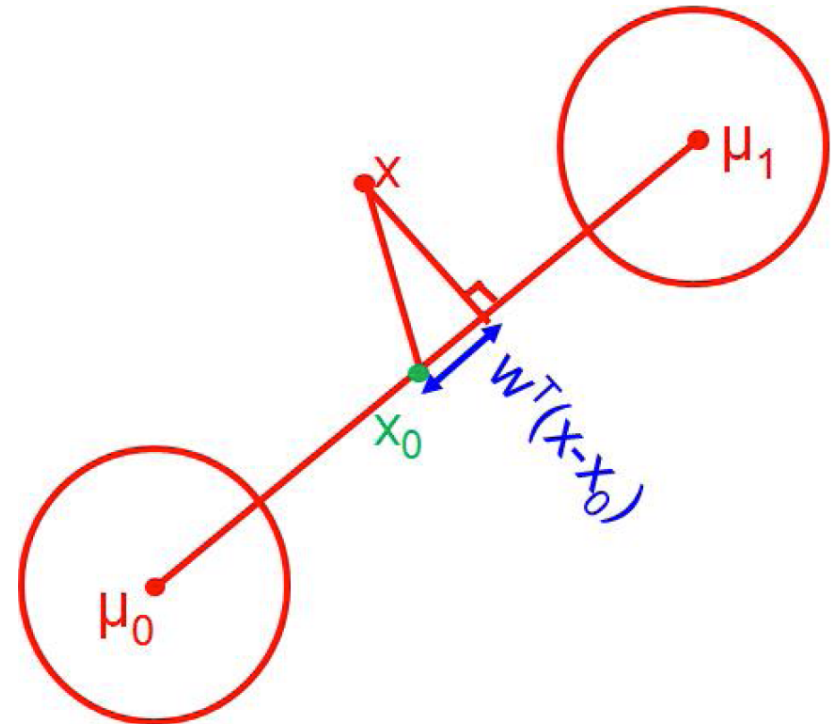


Binary Discriminant Analysis

Logistic Function



$$\text{sigm}(\eta) := \frac{1}{1 + \exp(-\eta)} = \frac{e^\eta}{e^\eta + 1}$$

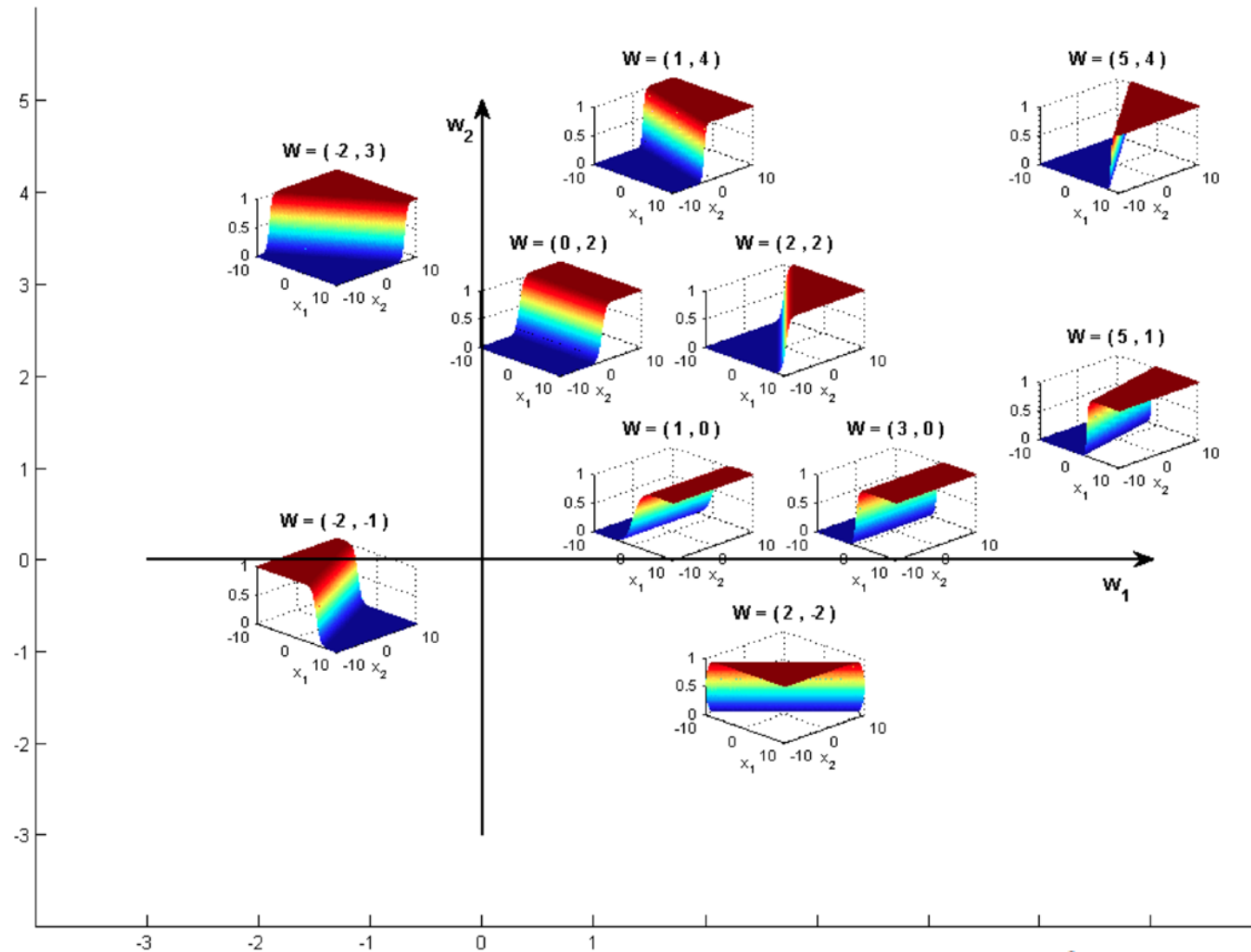


$$p(y = 1 | \mathbf{x}, \theta) = \sigma(\mathbf{w}^T (\mathbf{x} - \mathbf{x}_0))$$

$$\mathbf{w} = \beta_1 - \beta_0 = \Sigma^{-1}(\mu_1 - \mu_0)$$

$$\mathbf{x}_0 = \frac{1}{2}(\mu_1 + \mu_0) - (\mu_1 - \mu_0) \frac{\log(\pi_1/\pi_0)}{(\mu_1 - \mu_0)^T \Sigma^{-1}(\mu_1 - \mu_0)}$$

Logistic Regression

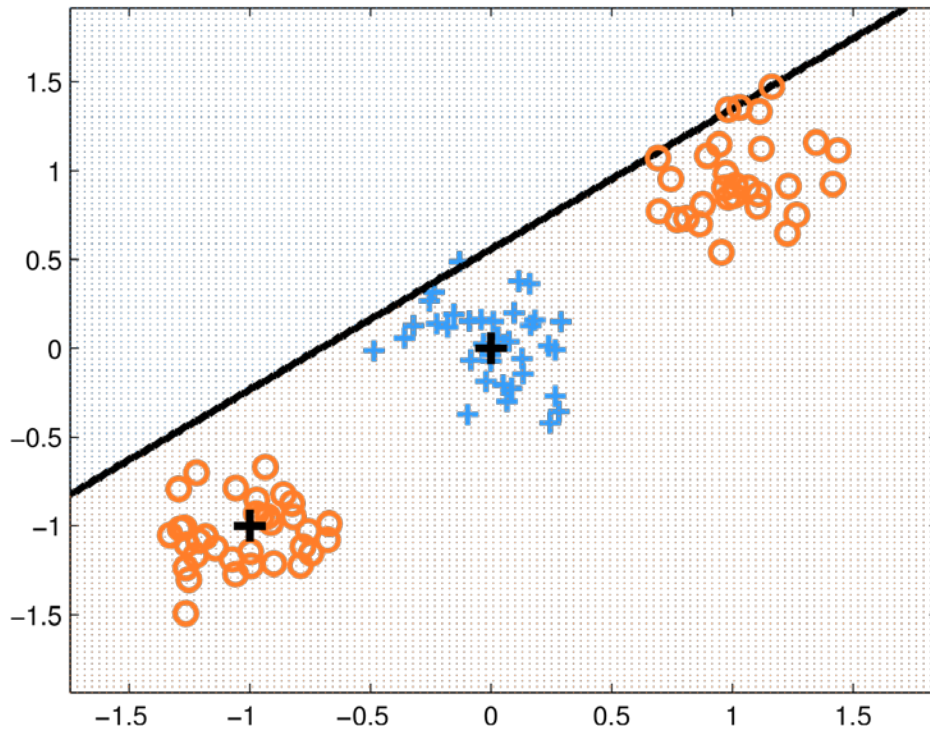


$$p(y|\mathbf{x}, \mathbf{w}) = \text{Ber}(y|\text{sigm}(\mathbf{w}^T \mathbf{x}))$$

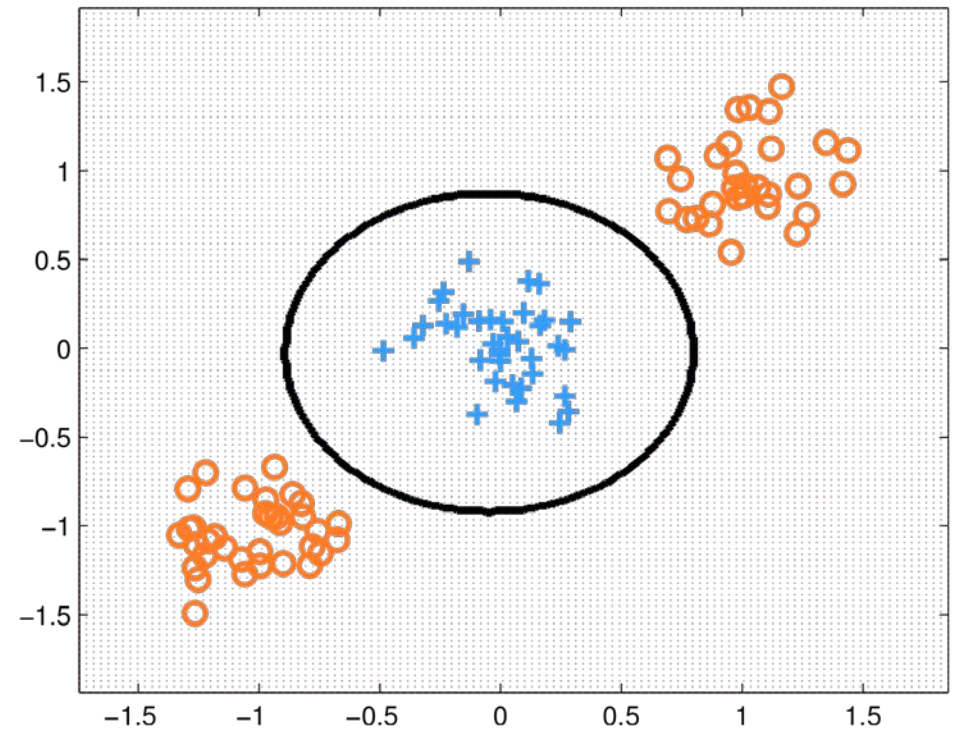
$$\text{sigm}(\eta) := \frac{1}{1 + \exp(-\eta)} = \frac{e^\eta}{e^\eta + 1}$$

Decision Boundaries

By assumption, a linear function of input features.

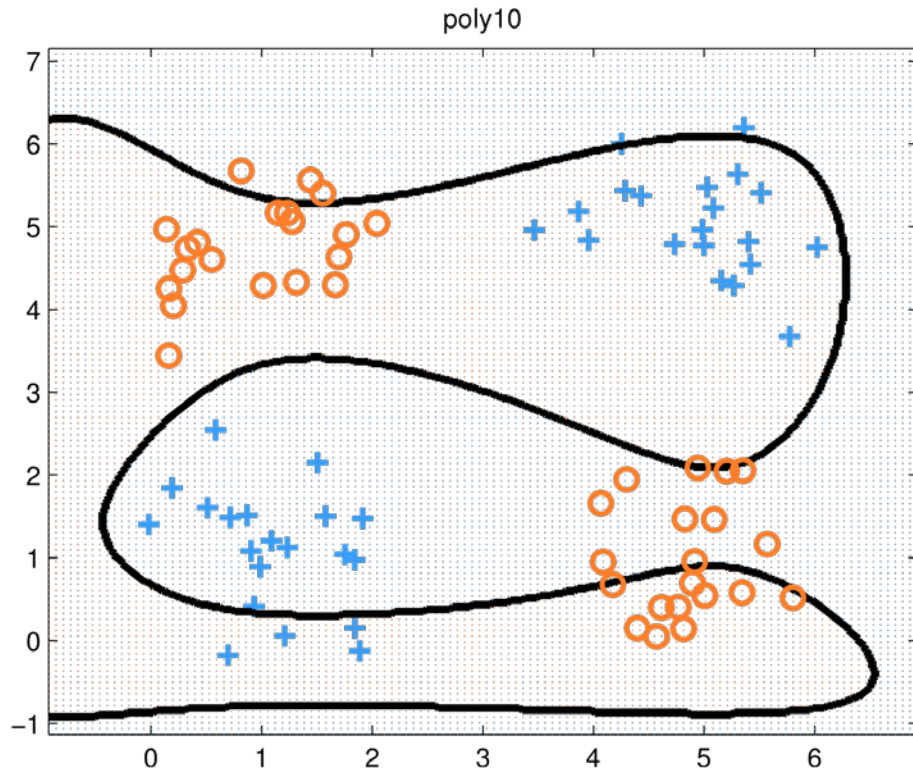


Raw features.

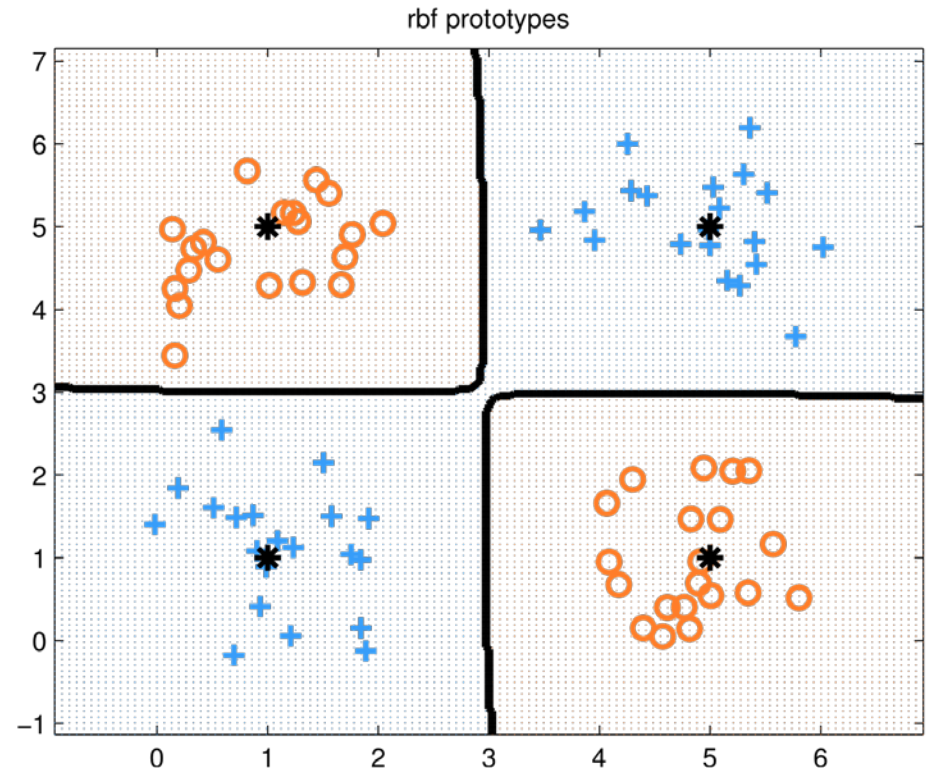


2nd-order polynomial expansion.

Importance of Features: XOR



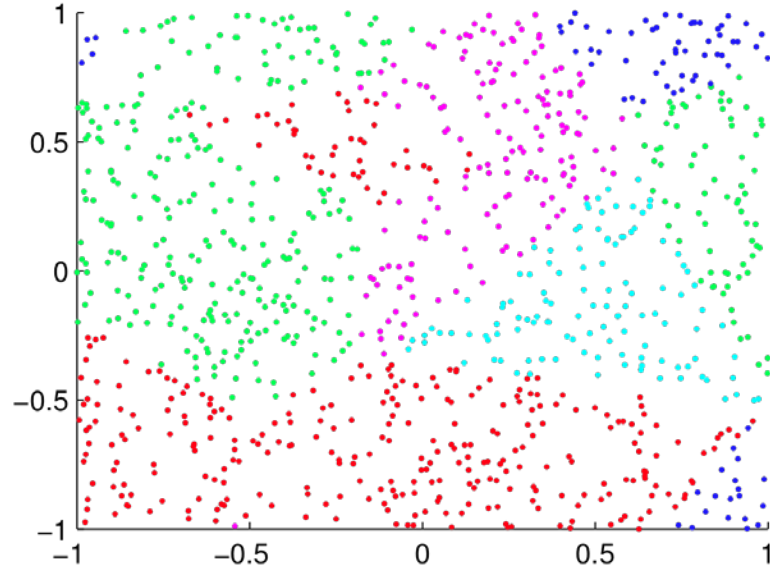
10th-order polynomial.



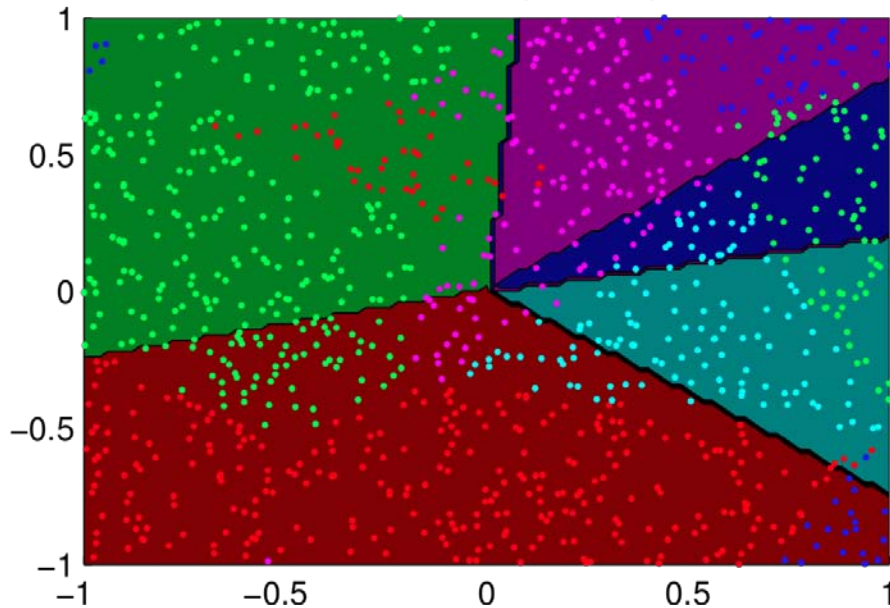
4 radial basis functions.

Avoid RBF placement via kernel methods...

Multinomial Logistic Regression



Linear Multinomial Logistic Regression



$$p(y|\mathbf{x}, \mathbf{W}) = \text{Cat}(y|\mathcal{S}(\mathbf{W}^T \mathbf{x}))$$

$$\mathcal{S}(\boldsymbol{\eta})_c = \frac{e^{\eta_c}}{\sum_{c'=1}^C e^{\eta_{c'}}}$$

as $T \rightarrow 0$

$$\mathcal{S}(\boldsymbol{\eta}/T)_c = \begin{cases} 1.0 & \text{if } c = \arg \max_{c'} \eta_{c'} \\ 0.0 & \text{otherwise} \end{cases}$$

Kernel-RBF Multinomial Logistic Regression

