



10-301/601 Introduction to Machine Learning

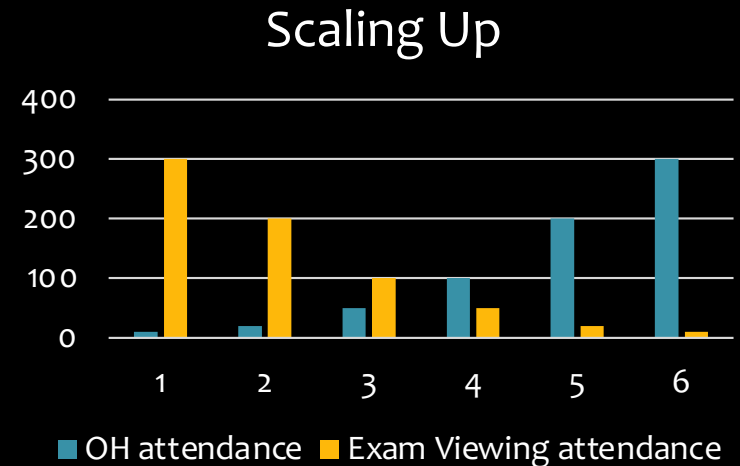
Machine Learning Department
School of Computer Science
Carnegie Mellon University

Neural Networks + Backpropagation

Matt Gormley
Lecture 12
Oct. 4, 2022

Reminders

- **Post-Exam Followup:**
 - Exam Viewing
 - Exit Poll: Exam 1
 - Grade Summary 1
- **Homework 4: Logistic Regression**
 - Out: Tue, Oct 4
 - Due: Thu, Oct 13 at 11:59pm
- **Homework 5: Neural Networks**
 - Out: Thu, Oct 13
 - Due: Thu, Oct 27 at 11:59pm



ARCHITECTURES

Neural Network Architectures

Even for a basic Neural Network, there are many design decisions to make:

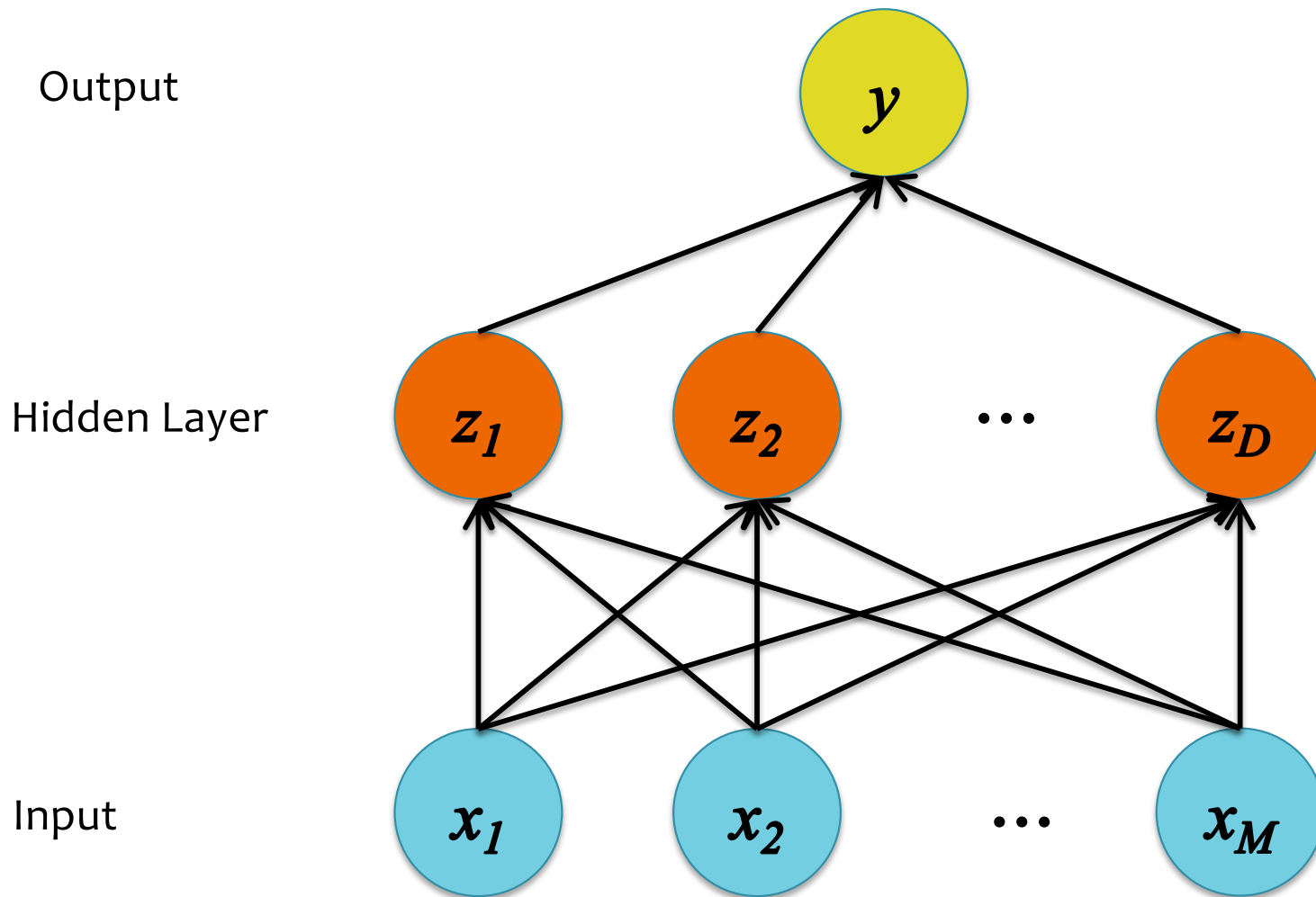
1. # of hidden layers (depth)
2. # of units per hidden layer (width)
3. Type of activation function (nonlinearity)
4. Form of objective function
5. How to initialize the parameters

BUILDING WIDER NETWORKS

$$D = M$$

Building a Neural Net

Q: How many hidden units, D , should we use?



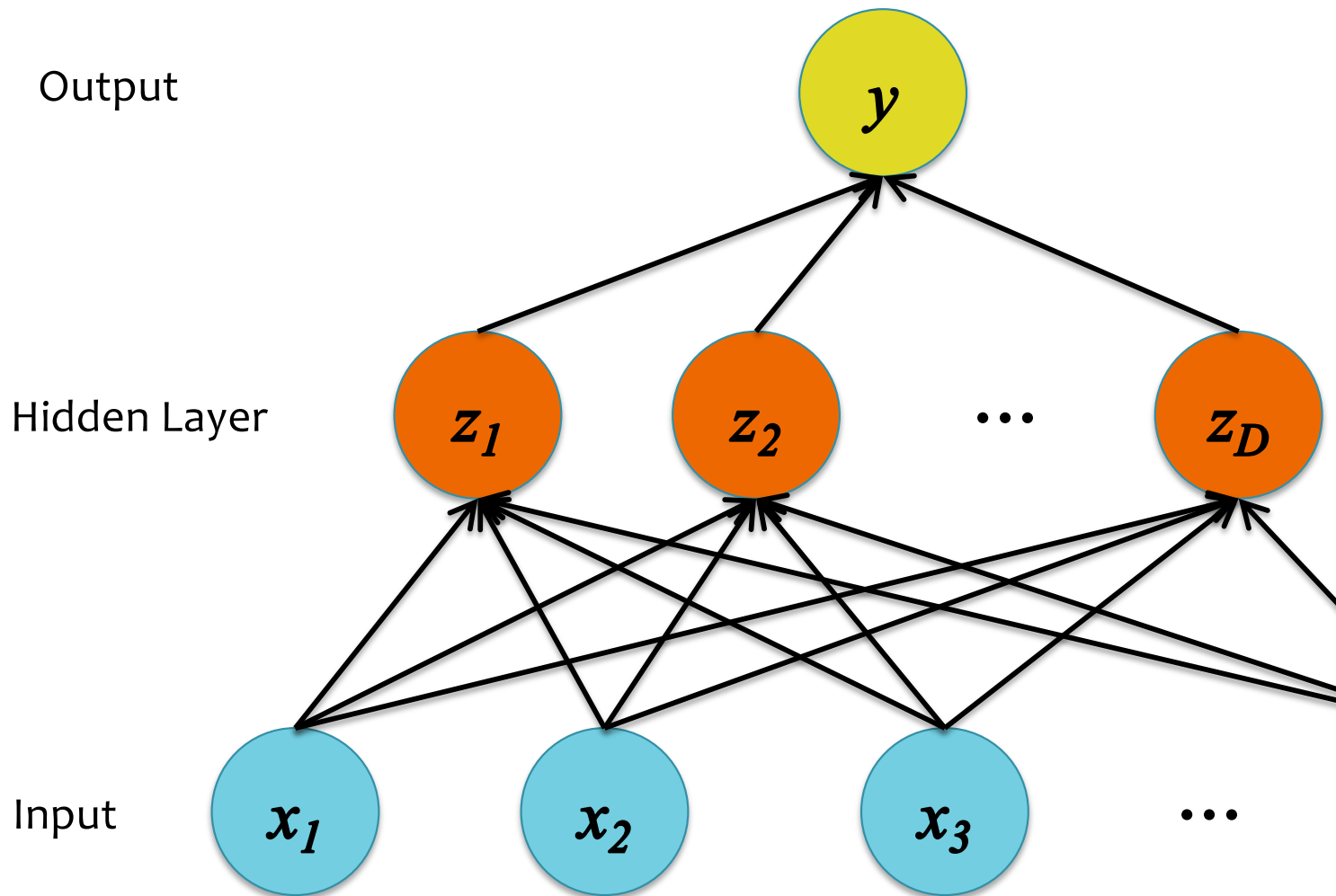
The hidden units could learn to be...

- a selection of the most useful features
- nonlinear combinations of the features
- a lower dimensional projection of the features
- a higher dimensional projection of the features
- a copy of the input features
- a mix of the above

$$D < M$$

Building a Neural Net

Q: How many hidden units, D , should we use?



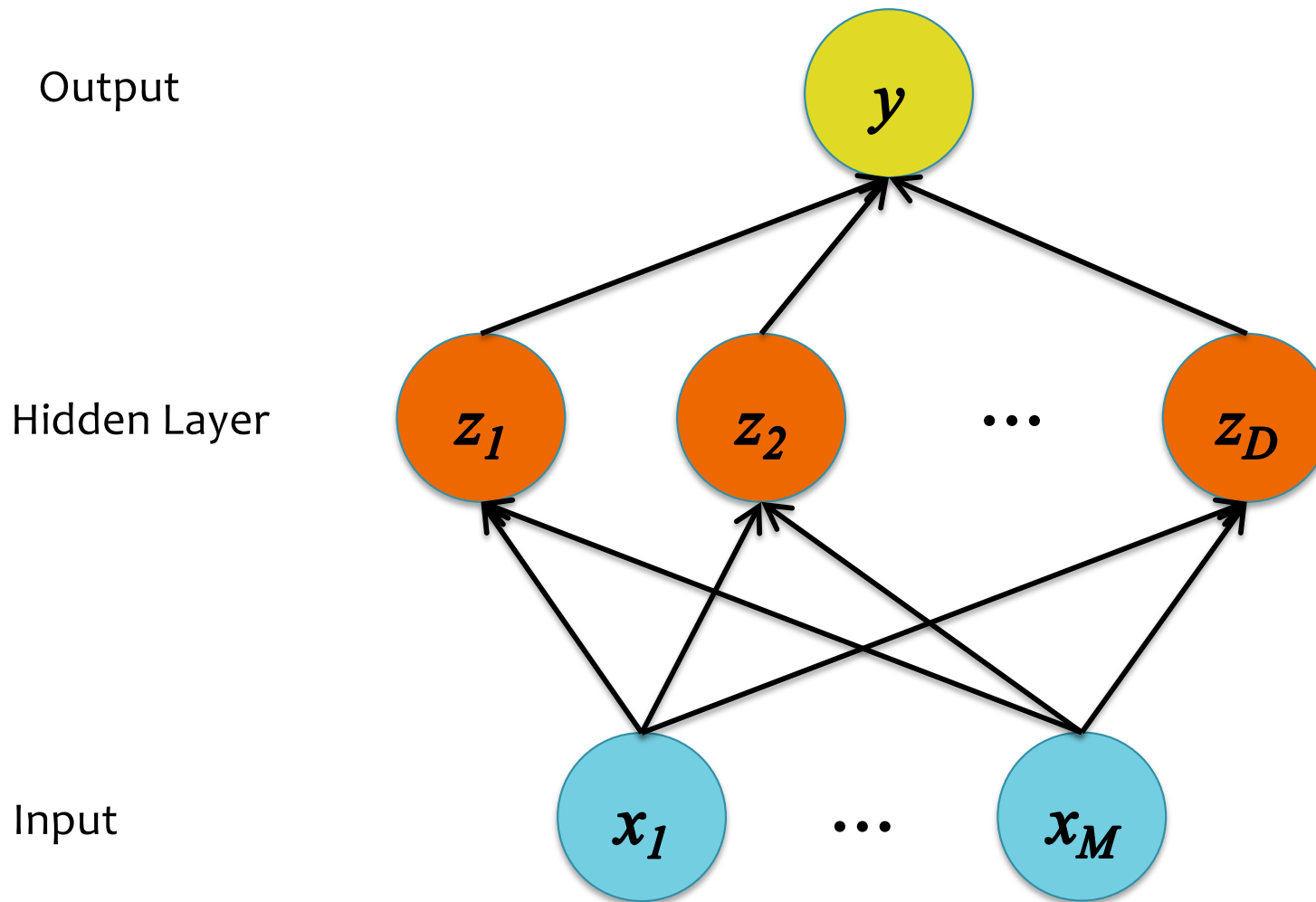
The hidden units could learn to be...

- a selection of the most useful features
- nonlinear combinations of the features
- a lower dimensional projection of the features
- a higher dimensional projection of the features
- a copy of the input features
- a mix of the above

$$D > M$$

Building a Neural Net

Q: How many hidden units, D , should we use?



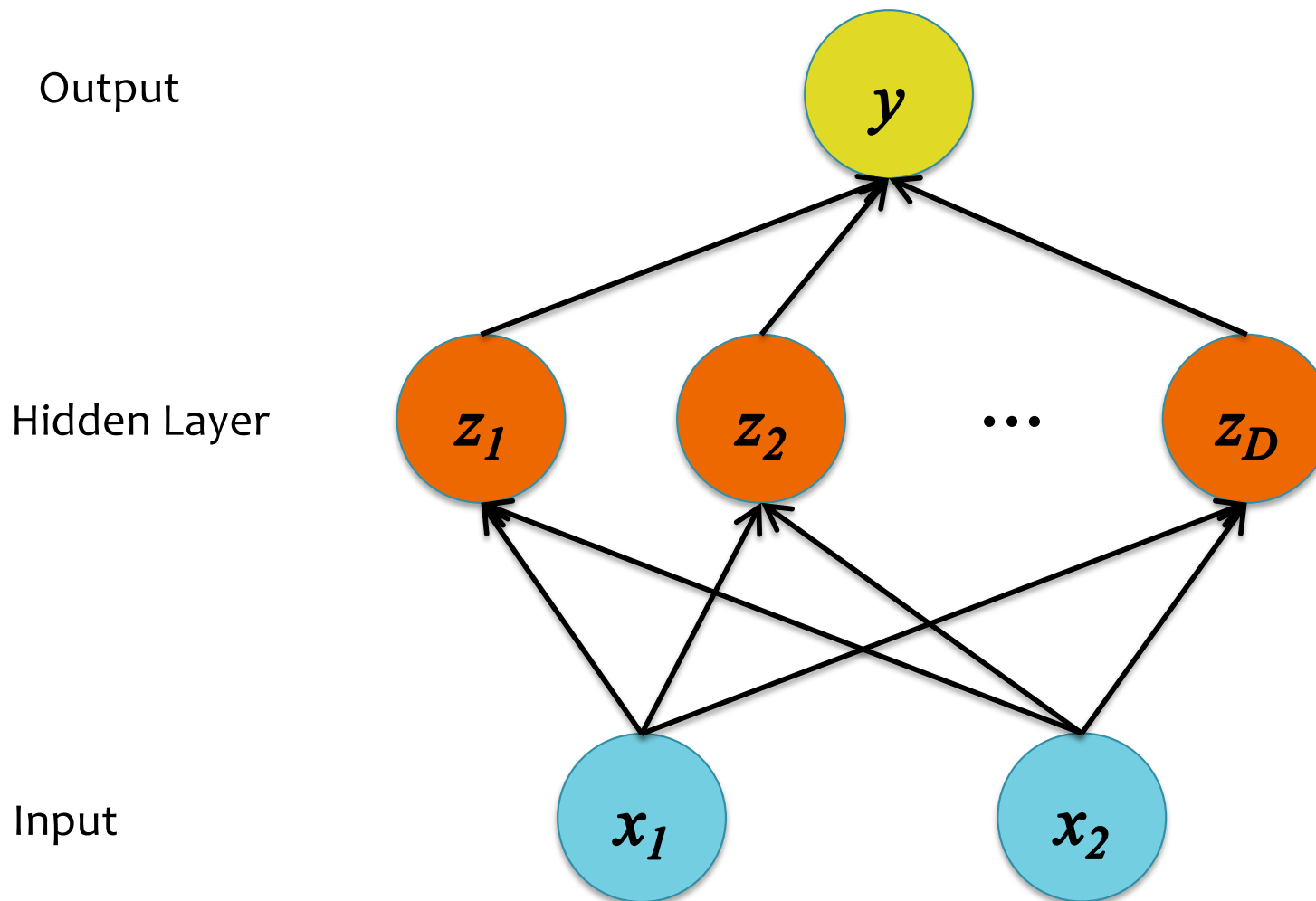
The hidden units could learn to be...

- a selection of the most useful features
- nonlinear combinations of the features
- a lower dimensional projection of the features
- a higher dimensional projection of the features
- a copy of the input features
- a mix of the above

$$D \geq M$$

Building a Neural Net

In the following examples, we have two input features, $M=2$, and we vary the number of hidden units, D .



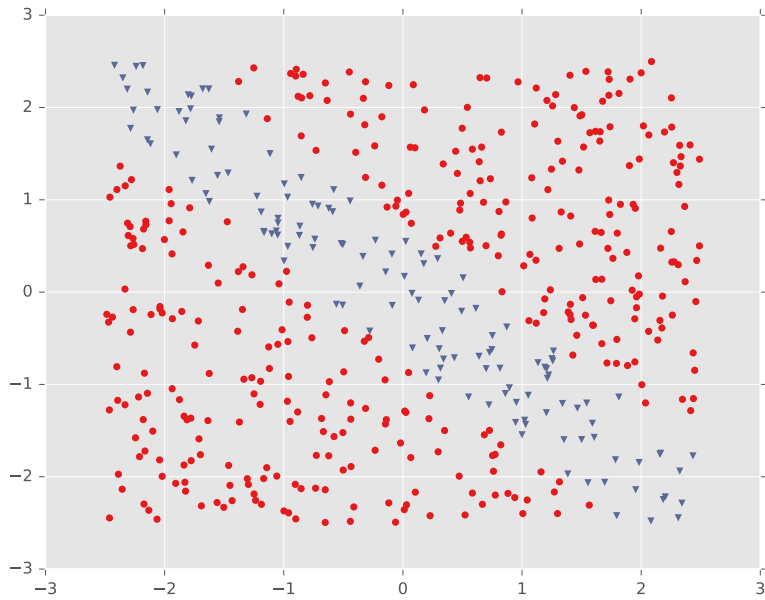
The hidden units could learn to be...

- a selection of the most useful features
- nonlinear combinations of the features
- a lower dimensional projection of the features
- a higher dimensional projection of the features
- a copy of the input features
- a mix of the above

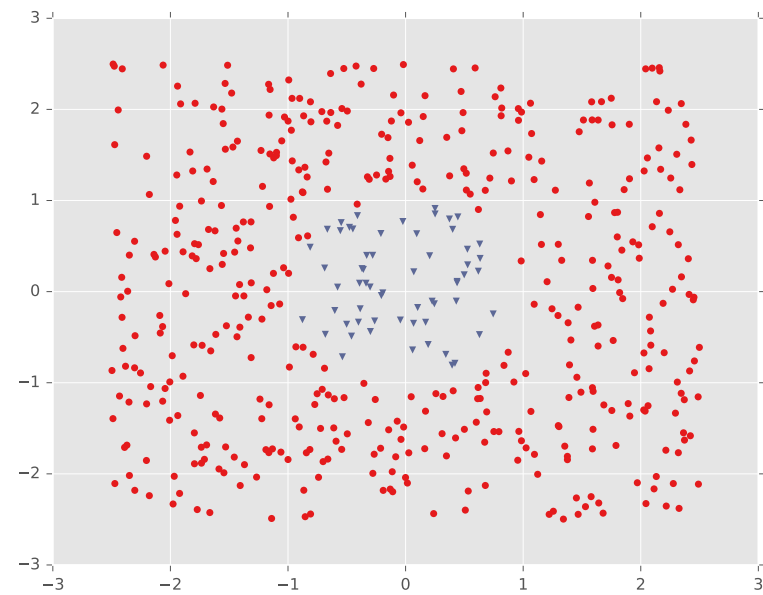
Examples 1 and 2

DECISION BOUNDARY EXAMPLES

Example #1: Diagonal Band



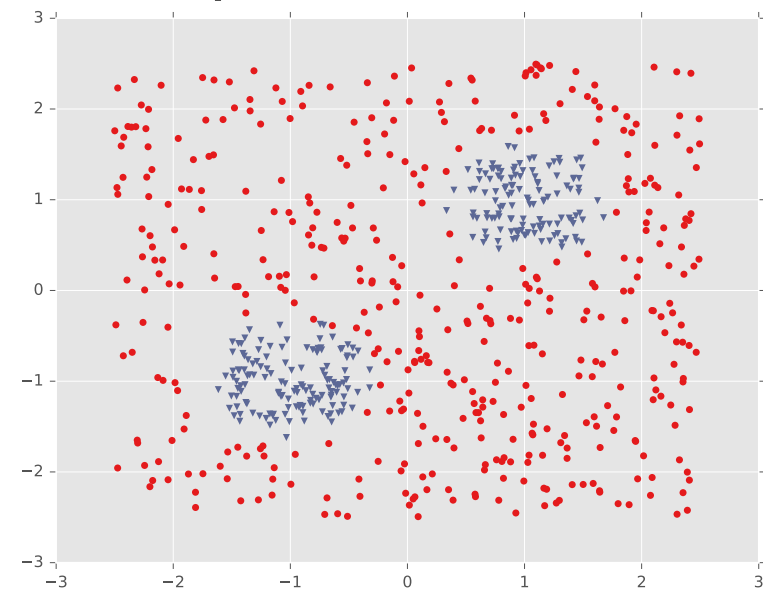
Example #2: One Pocket



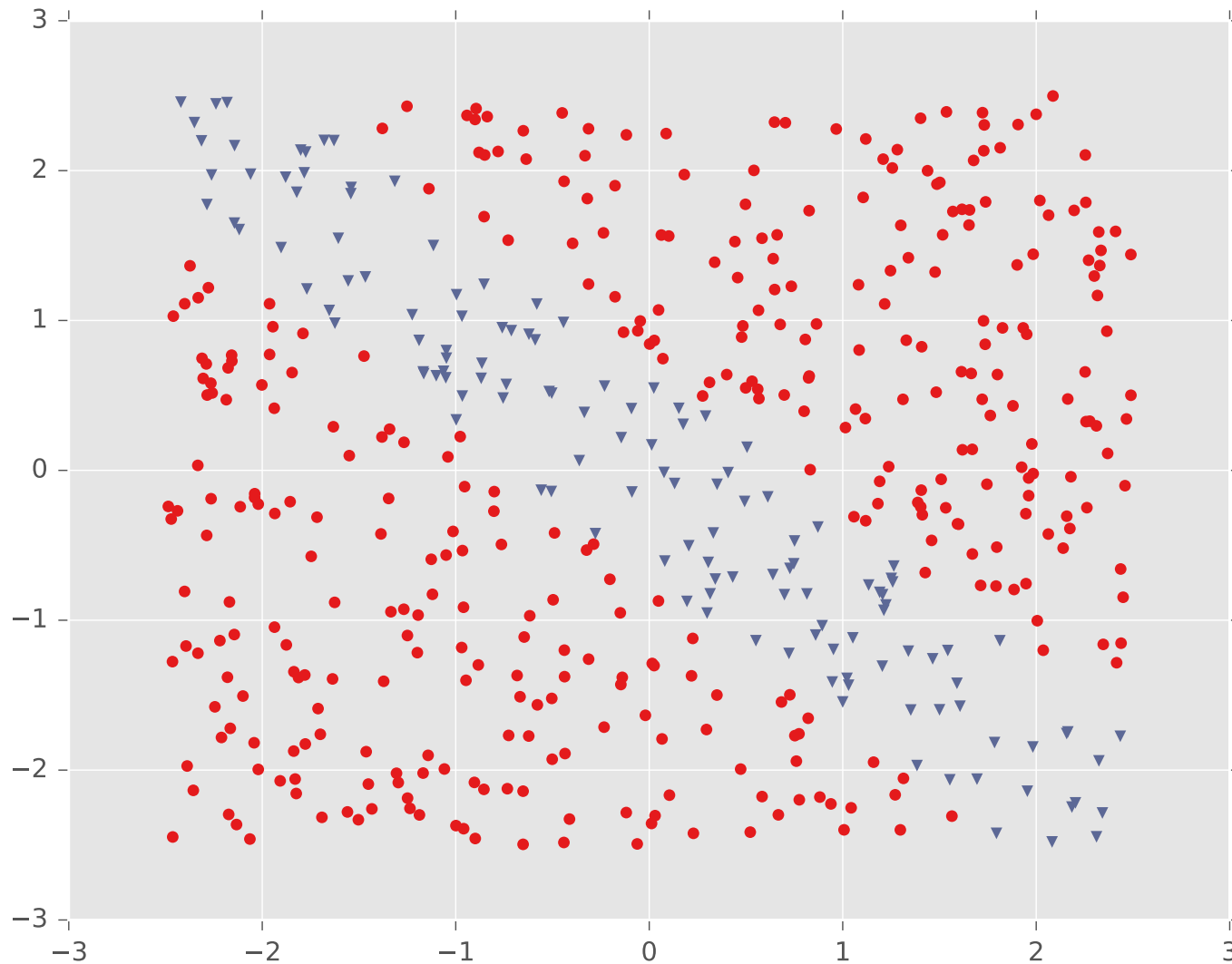
Example #3: Four Gaussians



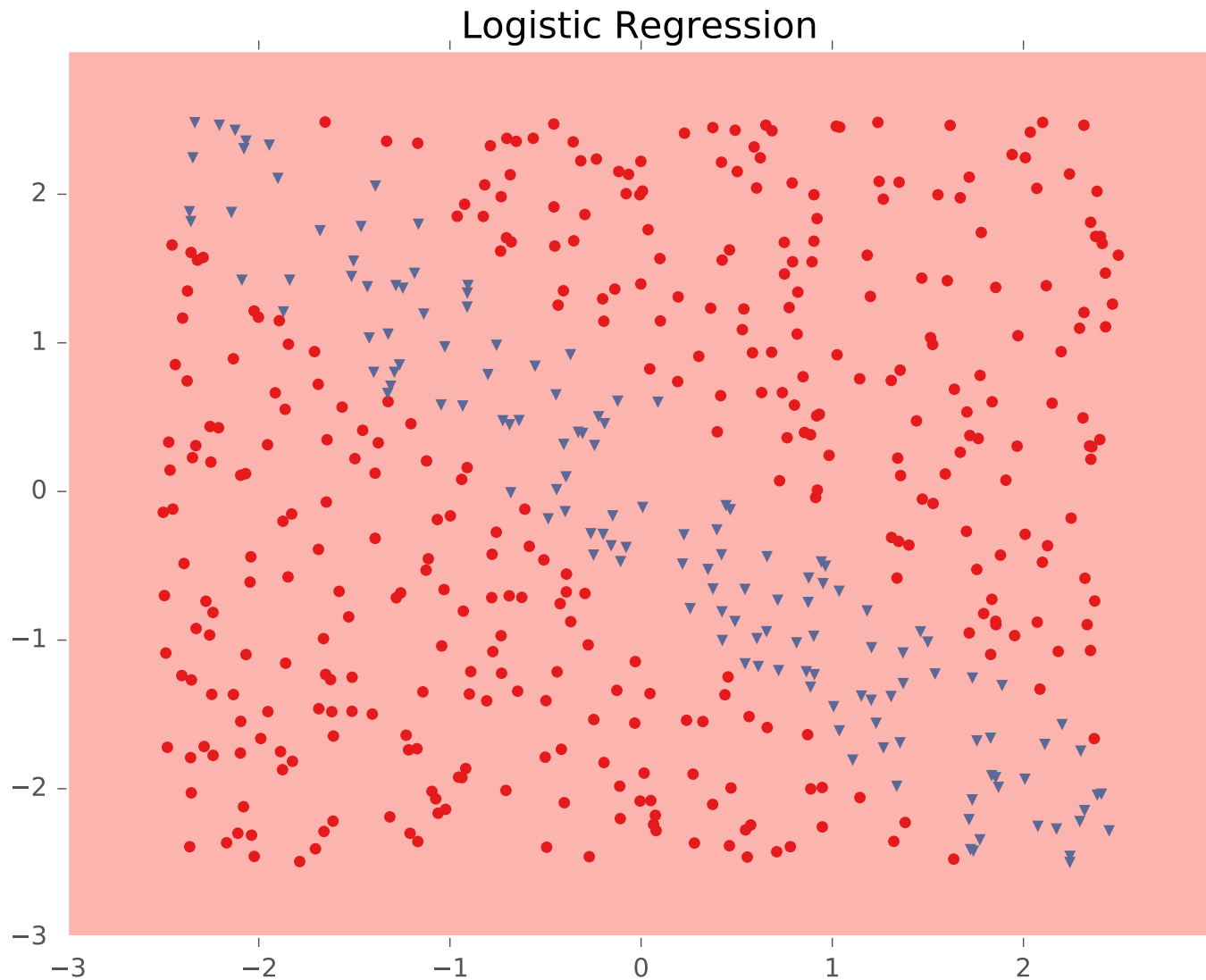
Example #4: Two Pockets



Example #1: Diagonal Band

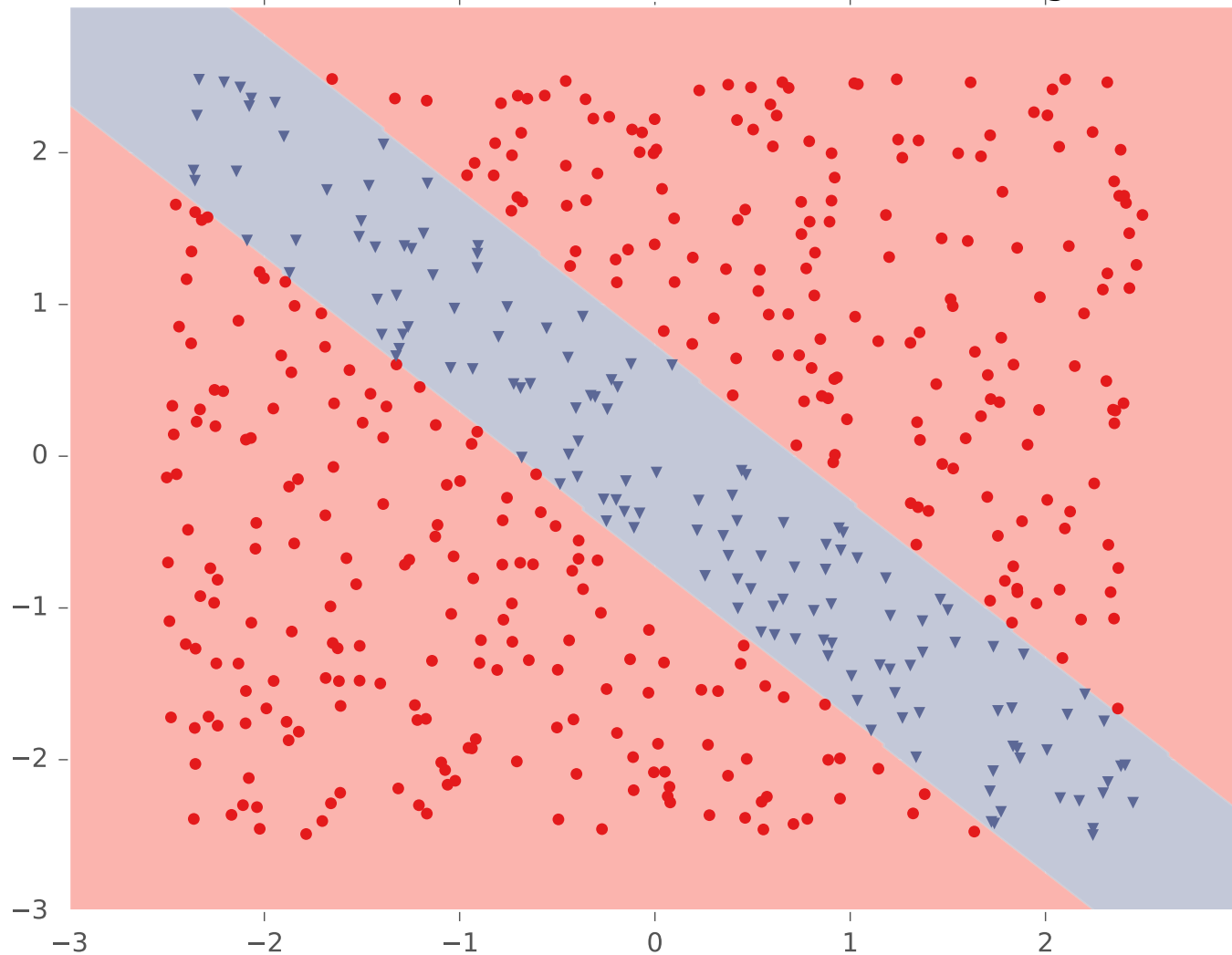


Example #1: Diagonal Band

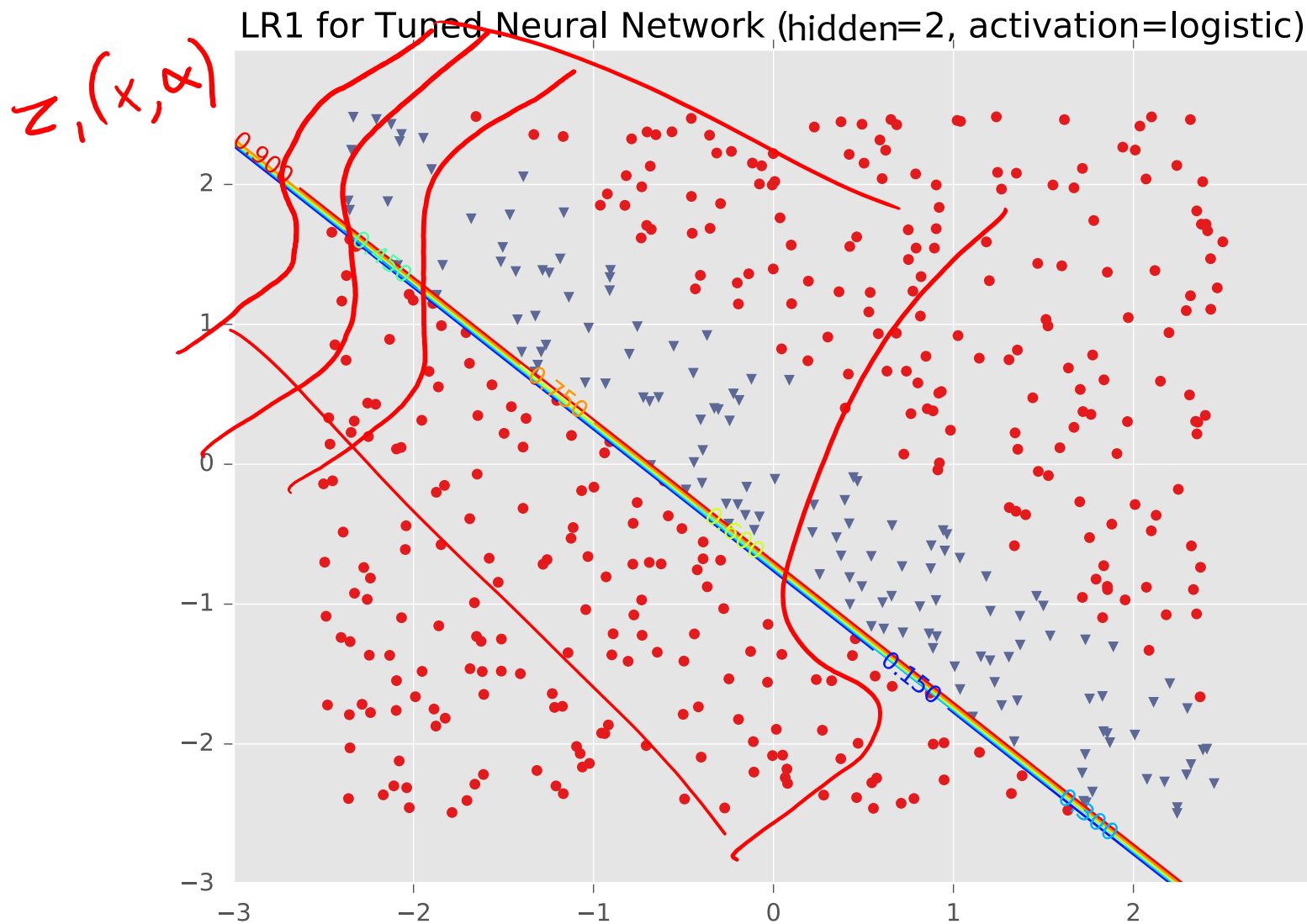


Example #1: Diagonal Band

Tuned Neural Network (hidden=2, activation=logistic)

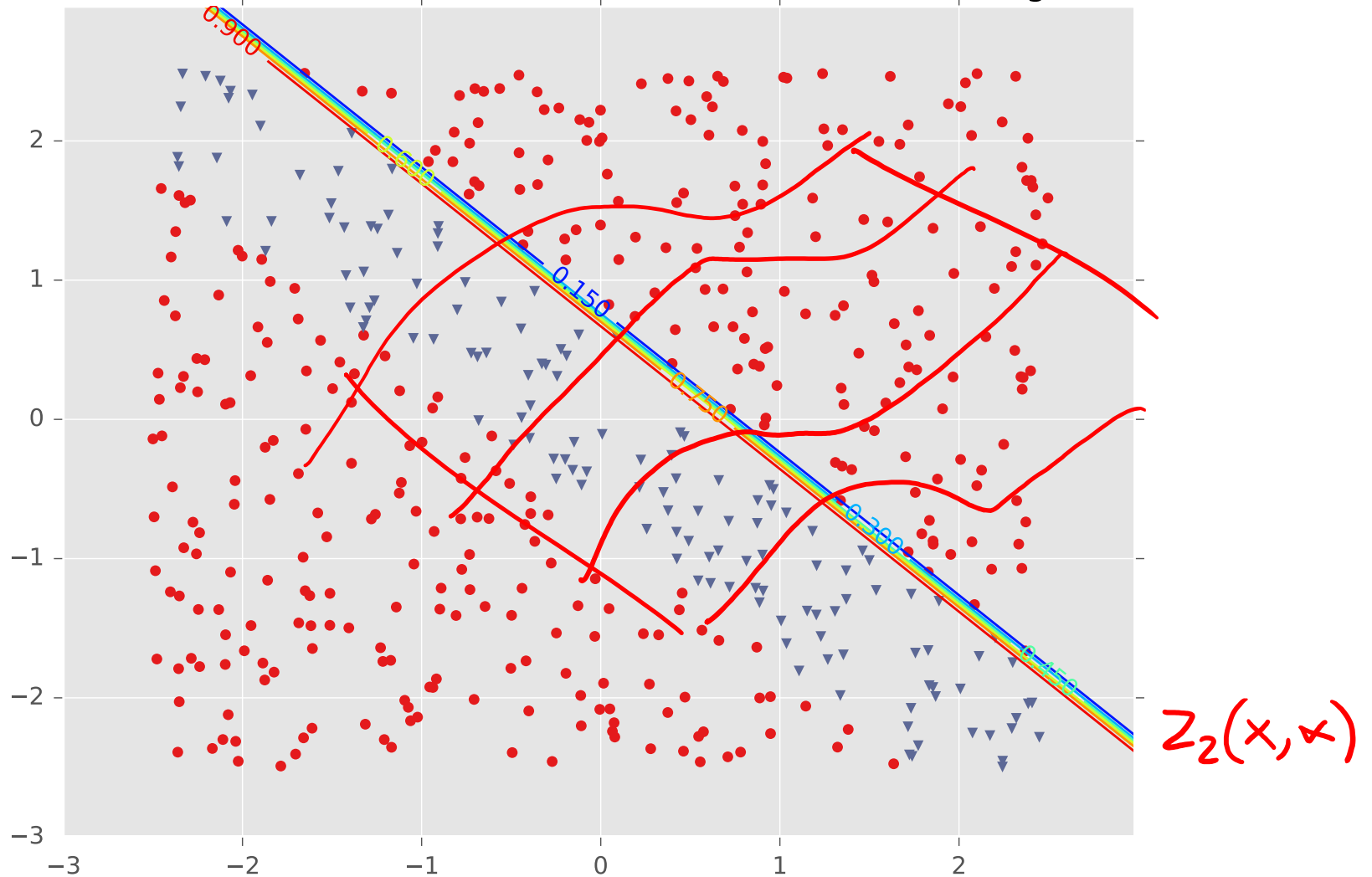


Example #1: Diagonal Band



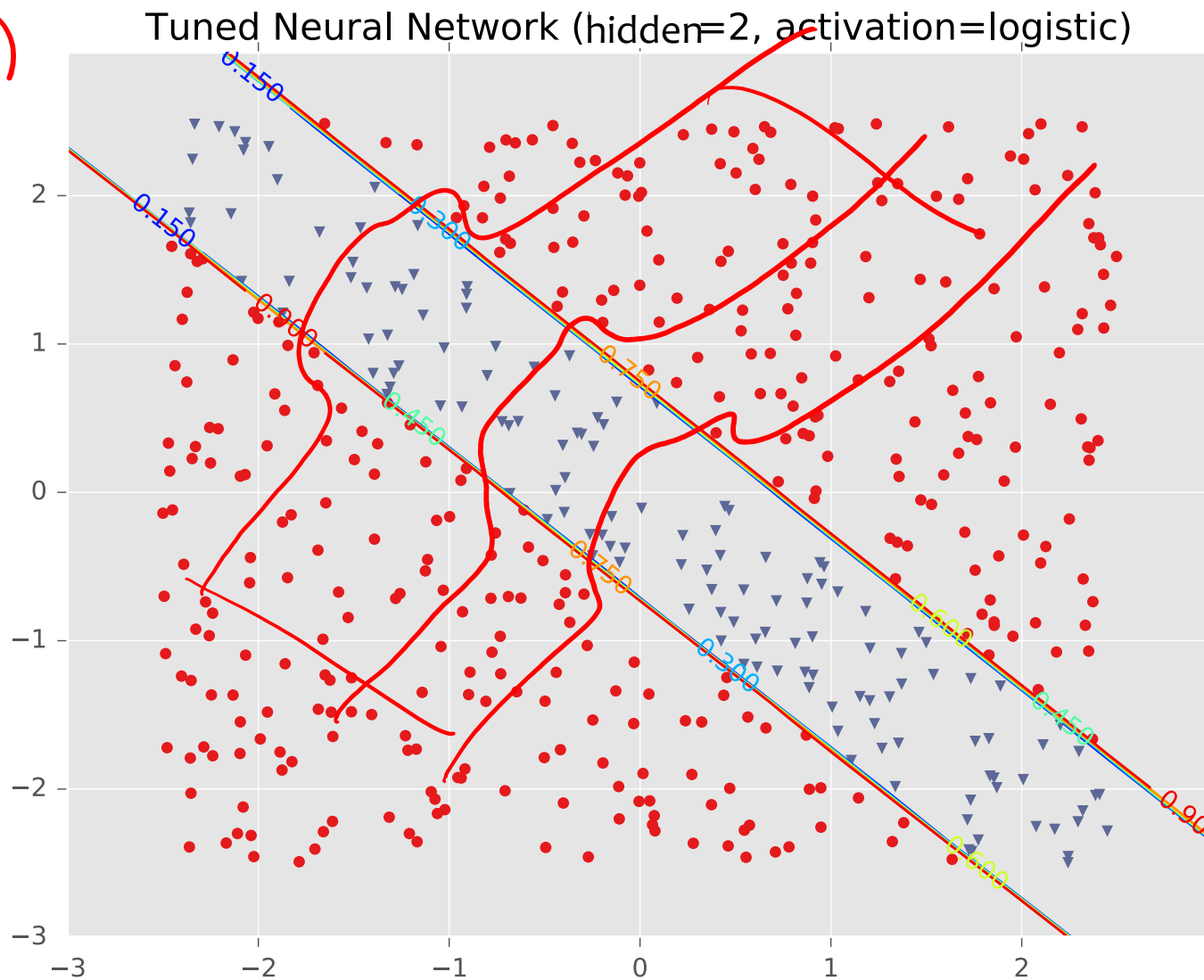
Example #1: Diagonal Band

LR2 for Tuned Neural Network (hidden=2, activation=logistic)



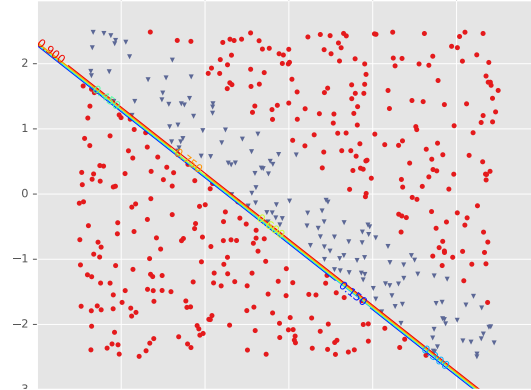
Example #1: Diagonal Band

$y(x, \alpha)$

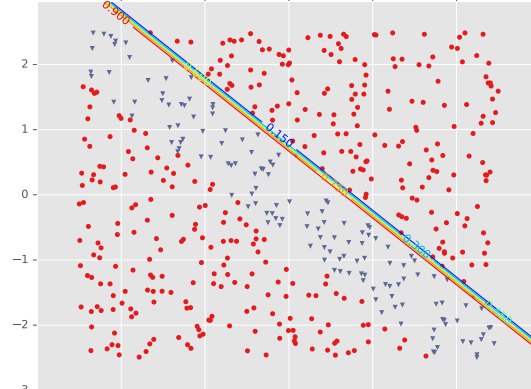


Example #1: Diagonal Band

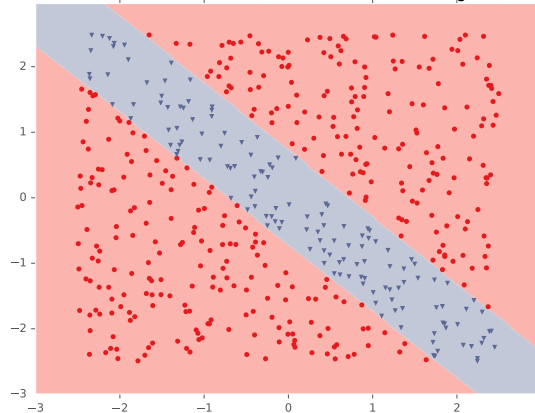
LR1 for Tuned Neural Network (hidden=2, activation=logistic)



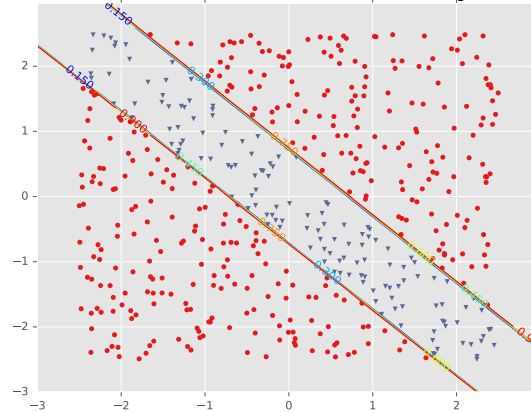
LR2 for Tuned Neural Network (hidden=2, activation=logistic)



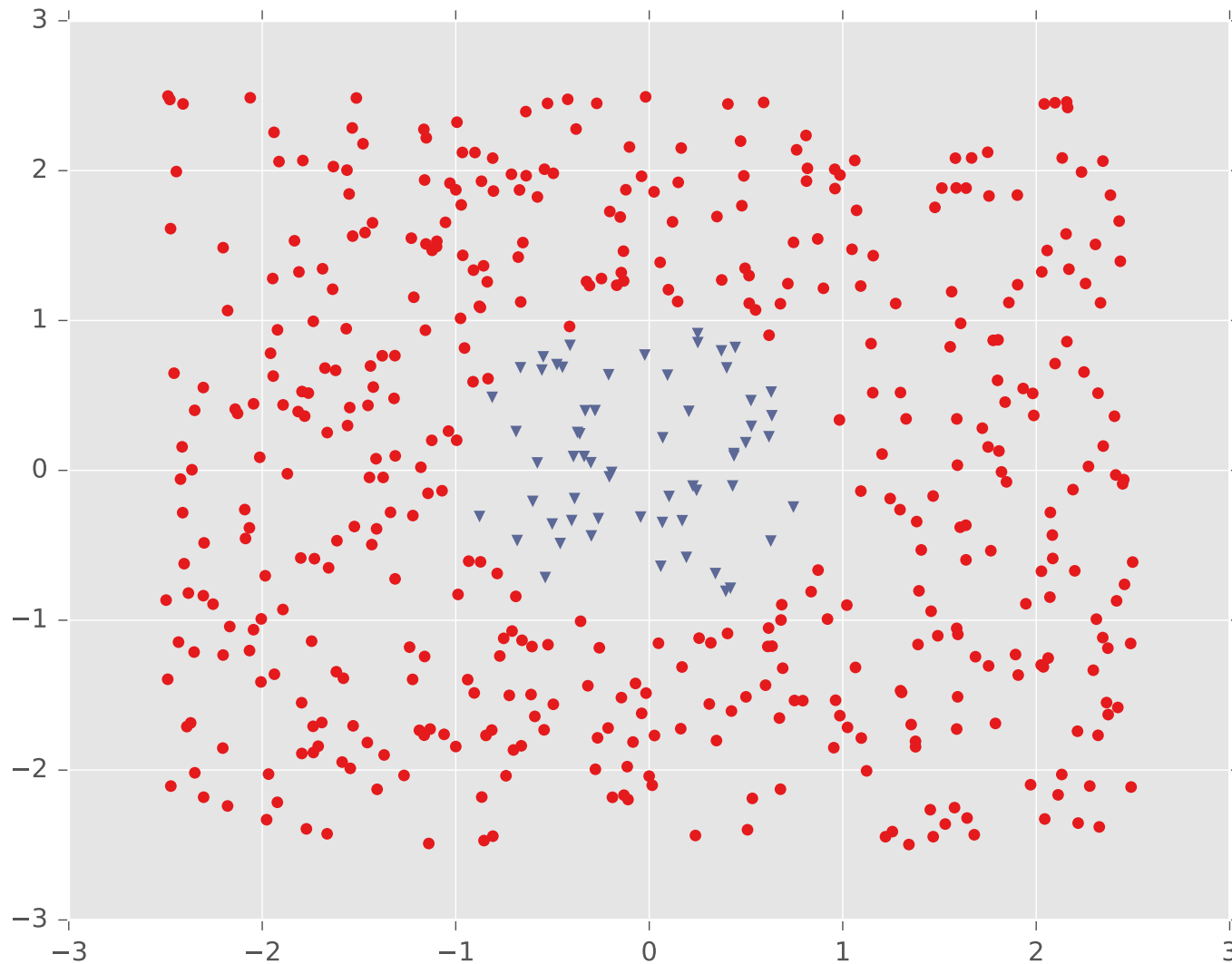
Tuned Neural Network (hidden=2, activation=logistic)



Tuned Neural Network (hidden=2, activation=logistic)



Example #2: One Pocket

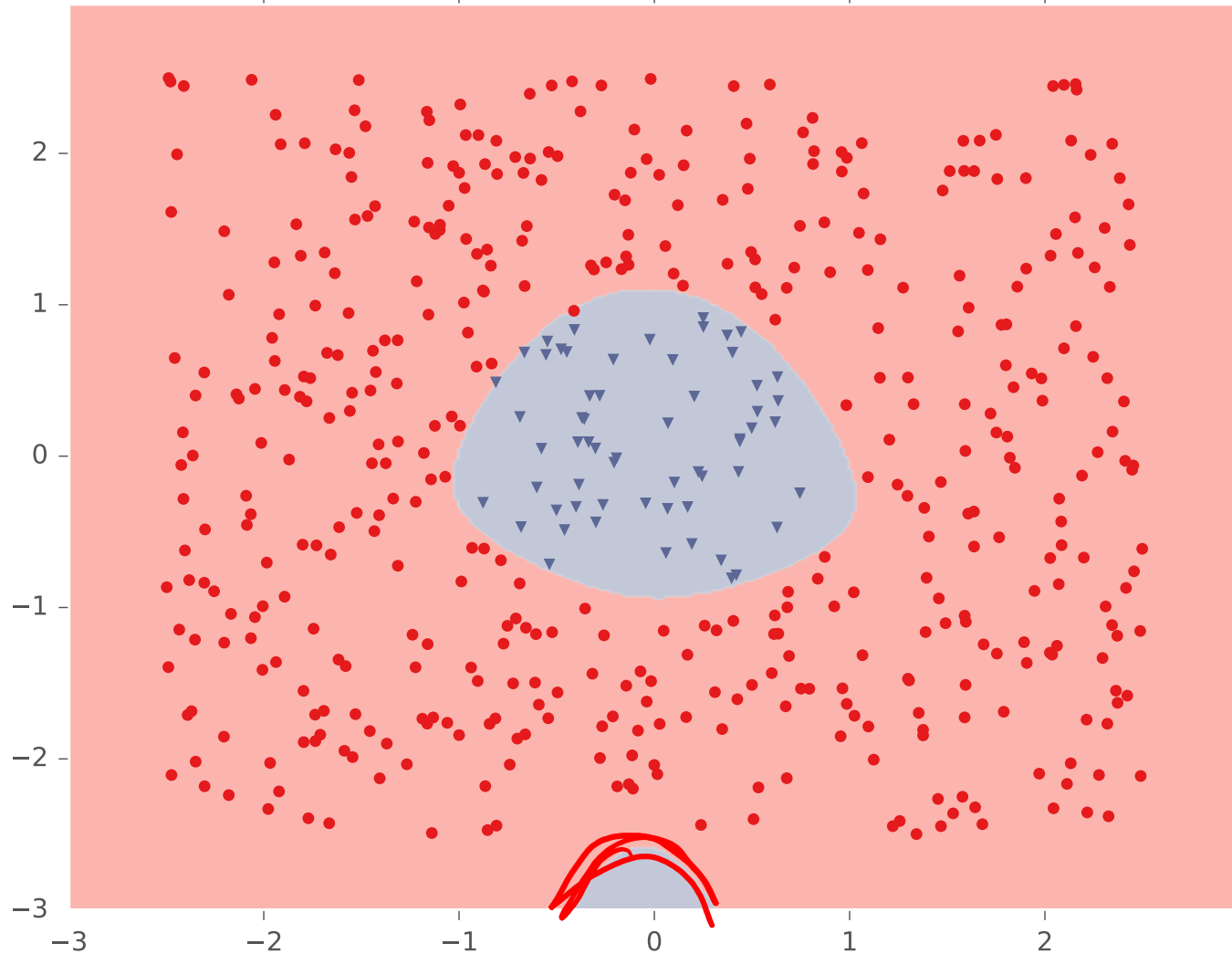


Example #2: One Pocket



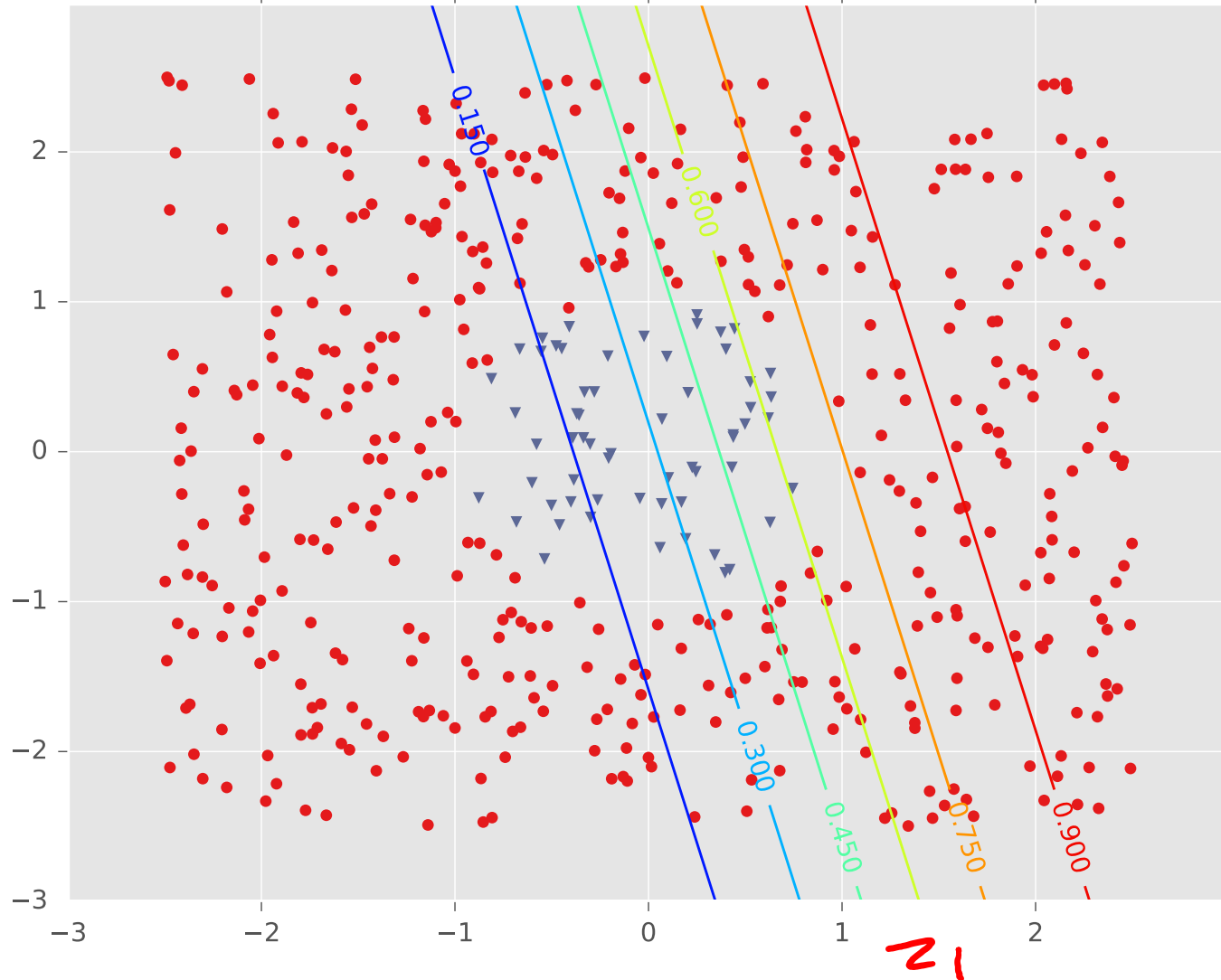
Example #2: One Pocket

Tuned Neural Network (hidden=3, activation=logistic)



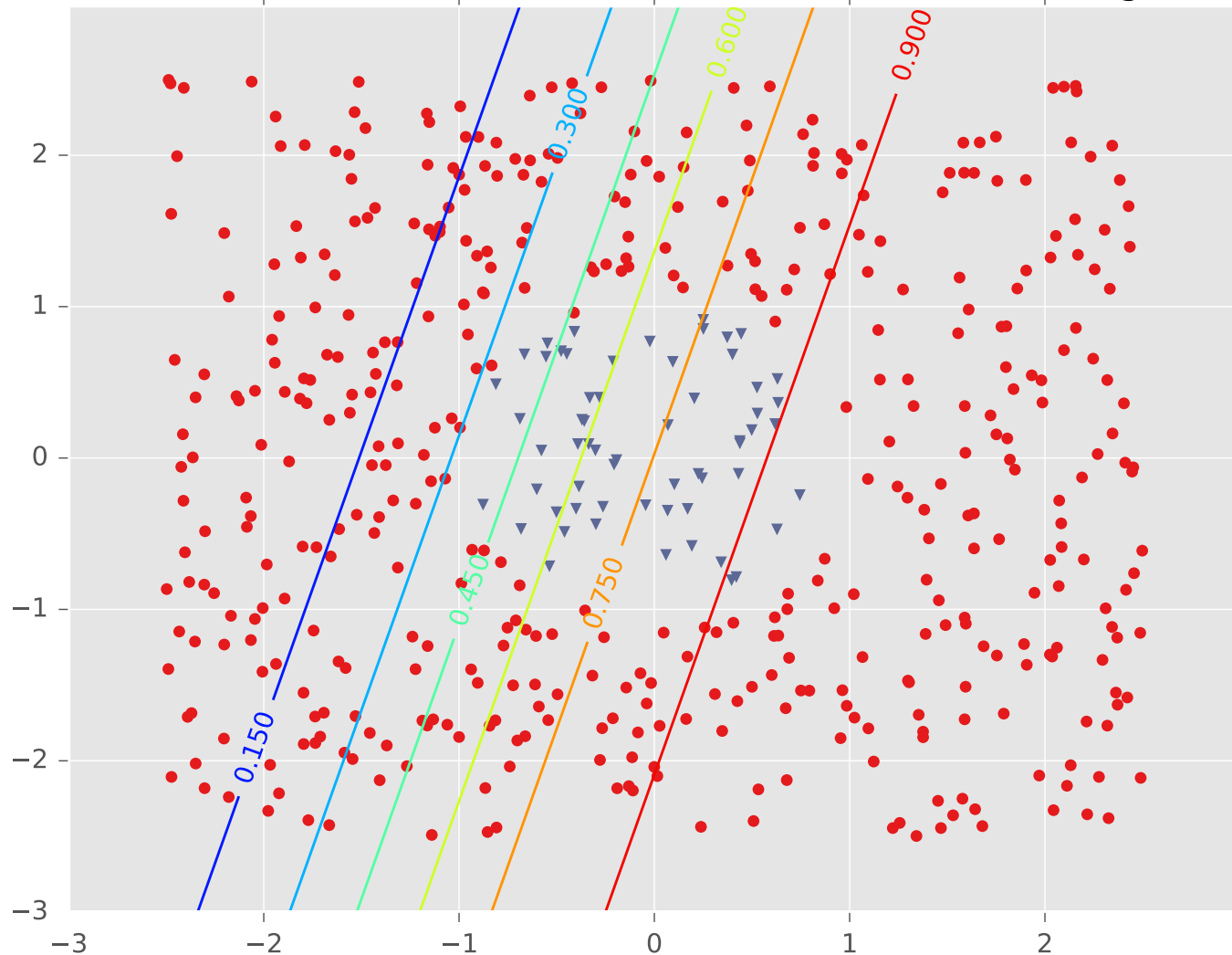
Example #2: One Pocket

LR1 for Tuned Neural Network (hidden=3, activation=logistic)



Example #2: One Pocket

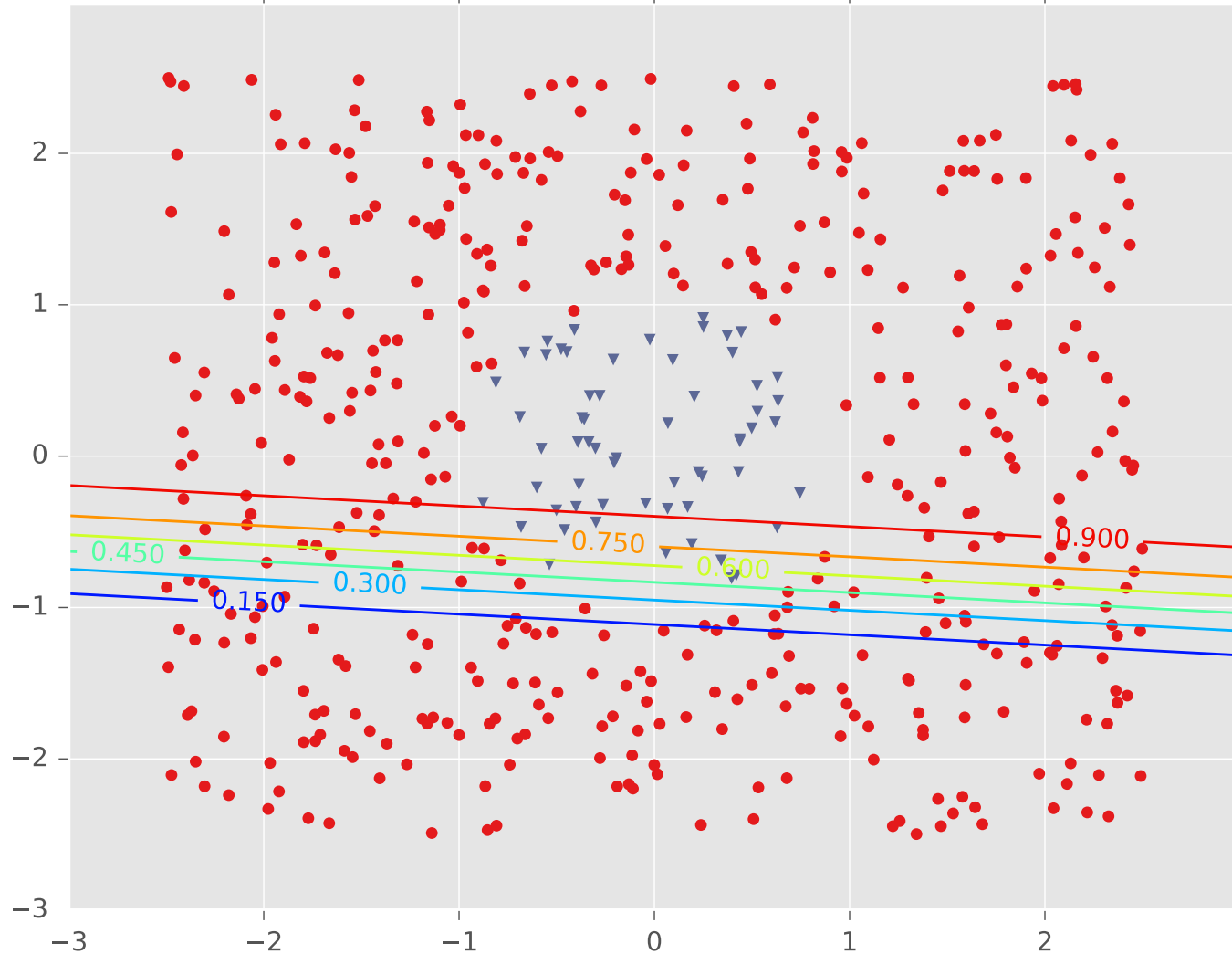
LR2 for Tuned Neural Network (hidden=3, activation=logistic)



z_2

Example #2: One Pocket

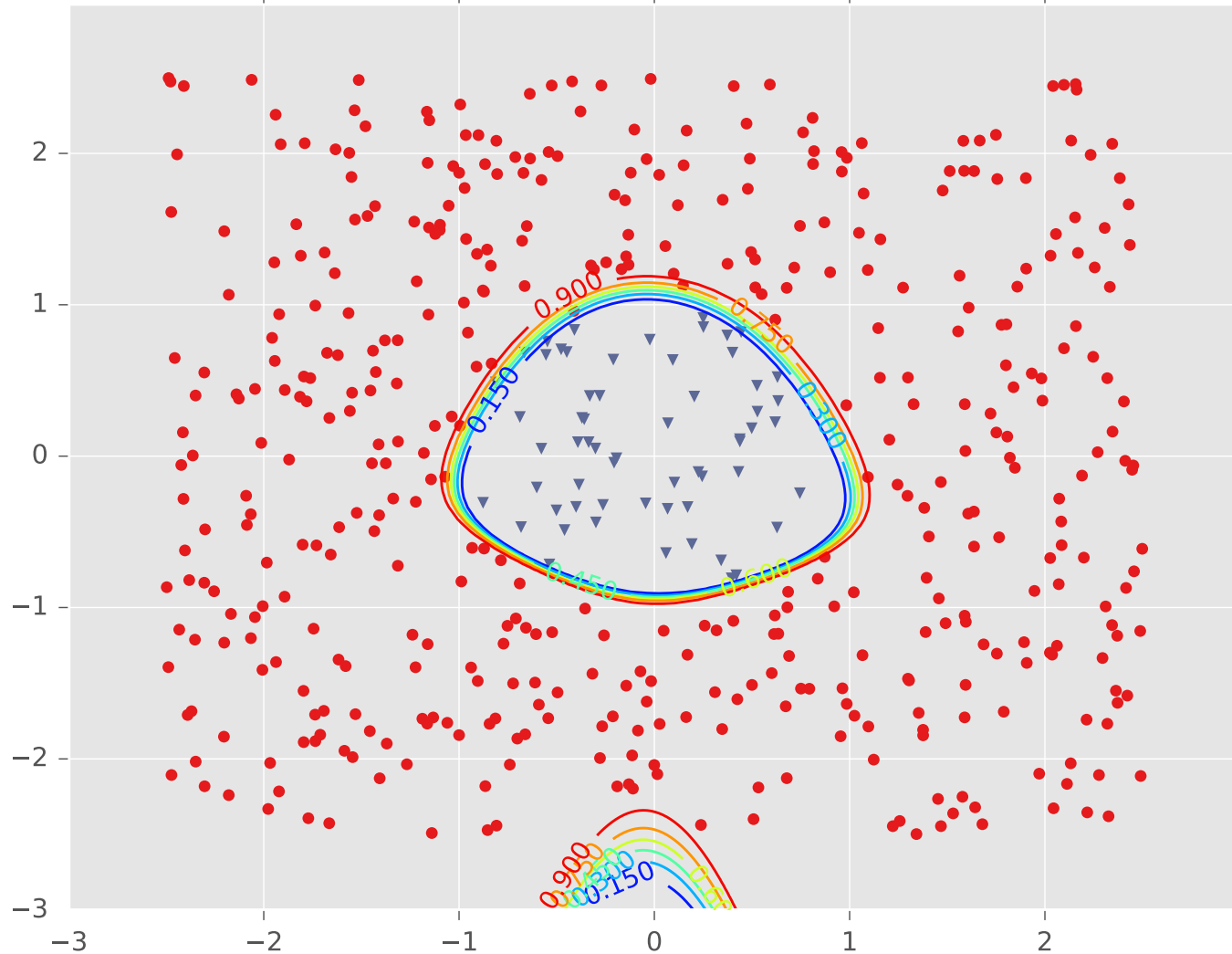
LR3 for Tuned Neural Network (hidden=3, activation=logistic)



z_3

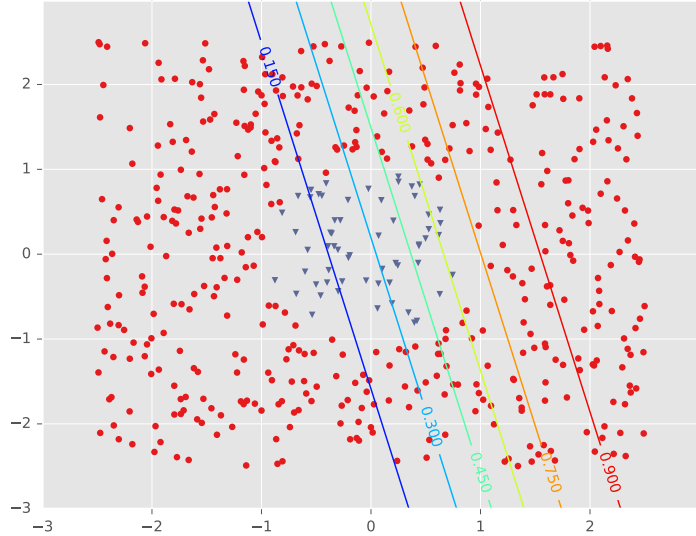
Example #2: One Pocket

Tuned Neural Network (hidden=3, activation=logistic)

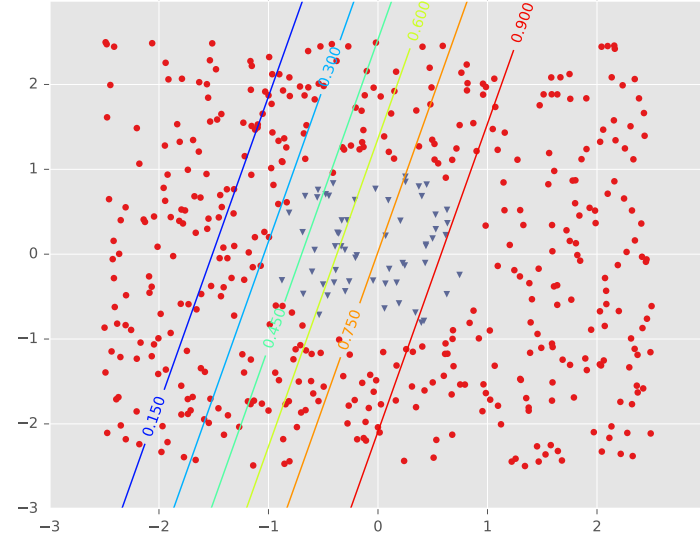


Example #2: One Pocket

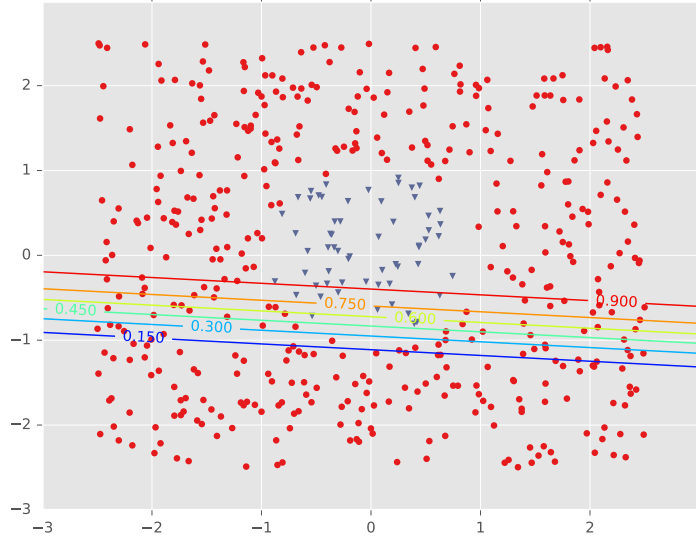
LR1 for Tuned Neural Network (hidden=3, activation=logistic)



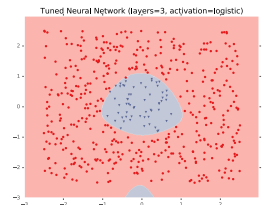
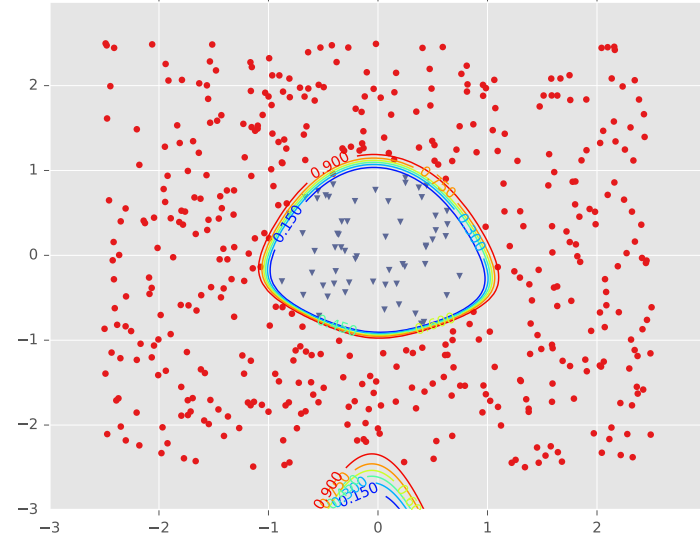
LR2 for Tuned Neural Network (hidden=3, activation=logistic)



LR3 for Tuned Neural Network (hidden=3, activation=logistic)



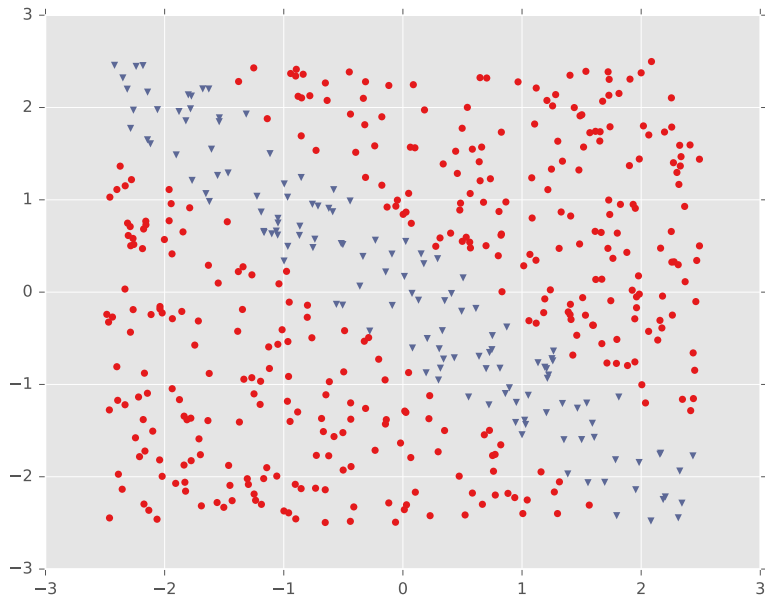
Tuned Neural Network (hidden=3, activation=logistic)



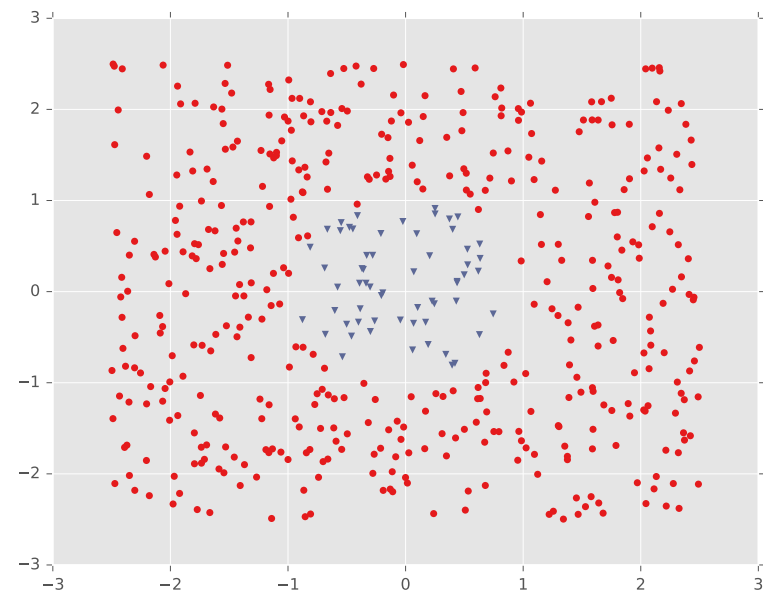
Examples 3 and 4

DECISION BOUNDARY EXAMPLES

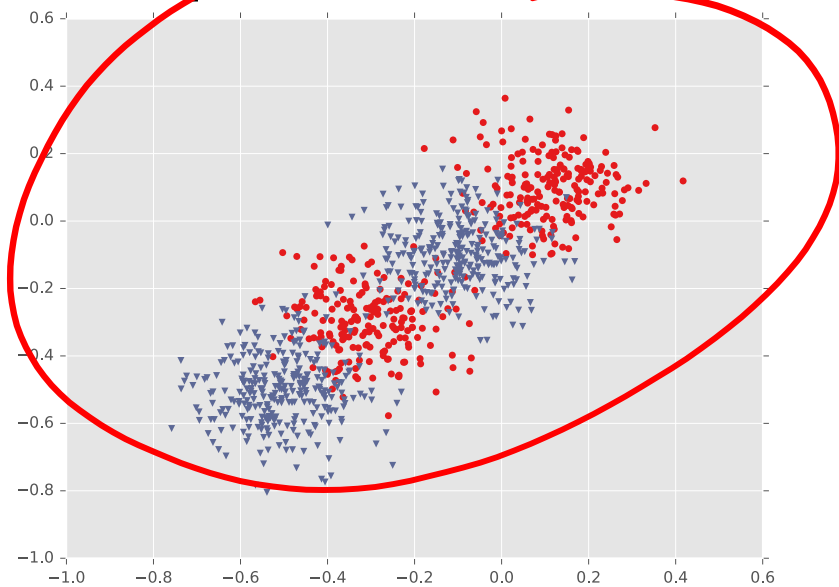
Example #1: Diagonal Band



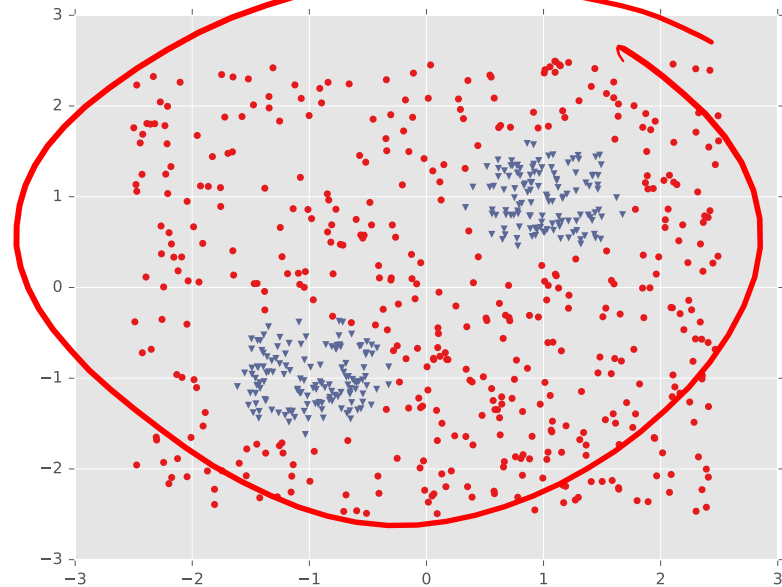
Example #2: One Pocket



Example #3: Four Gaussians



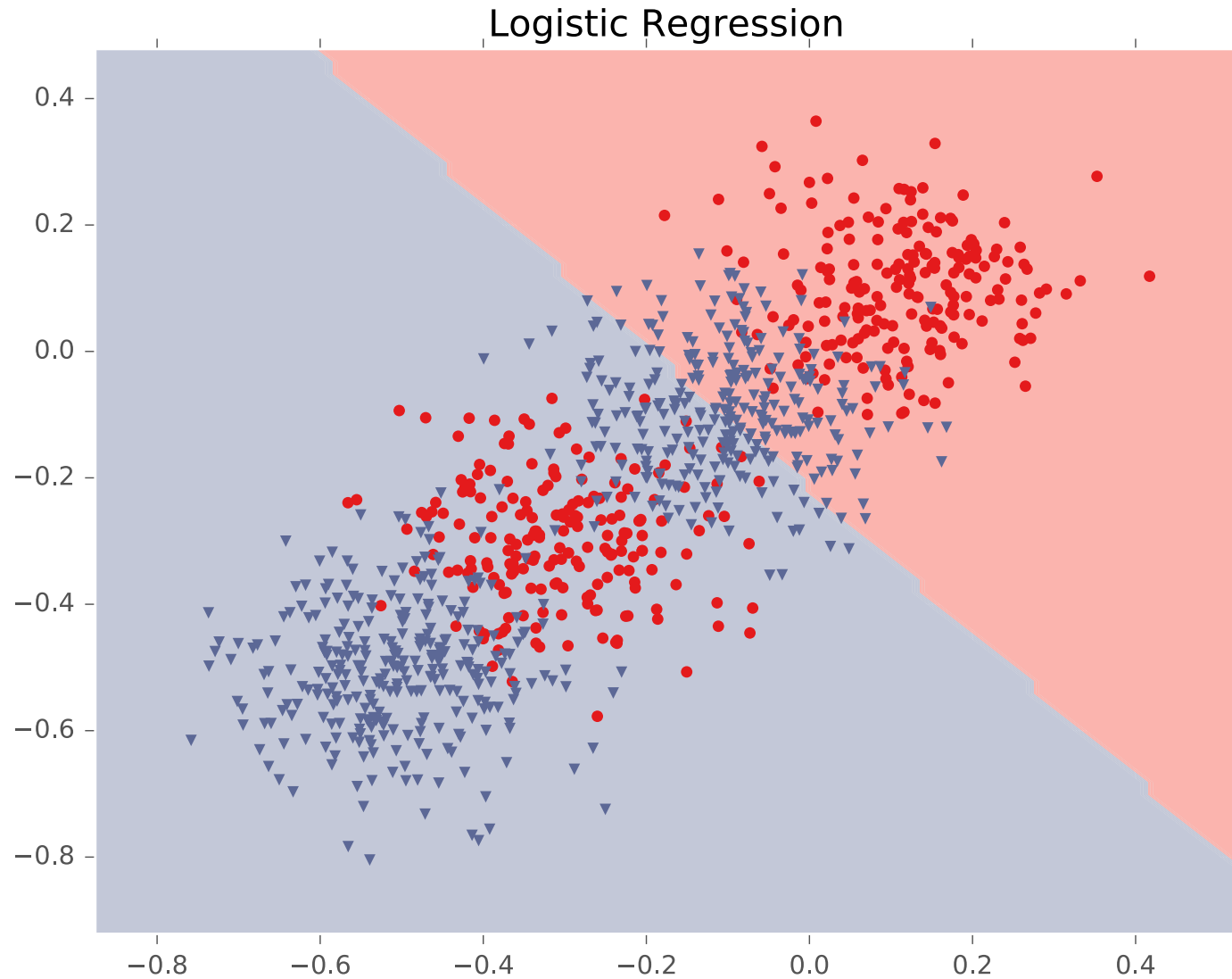
Example #4: Two Pockets



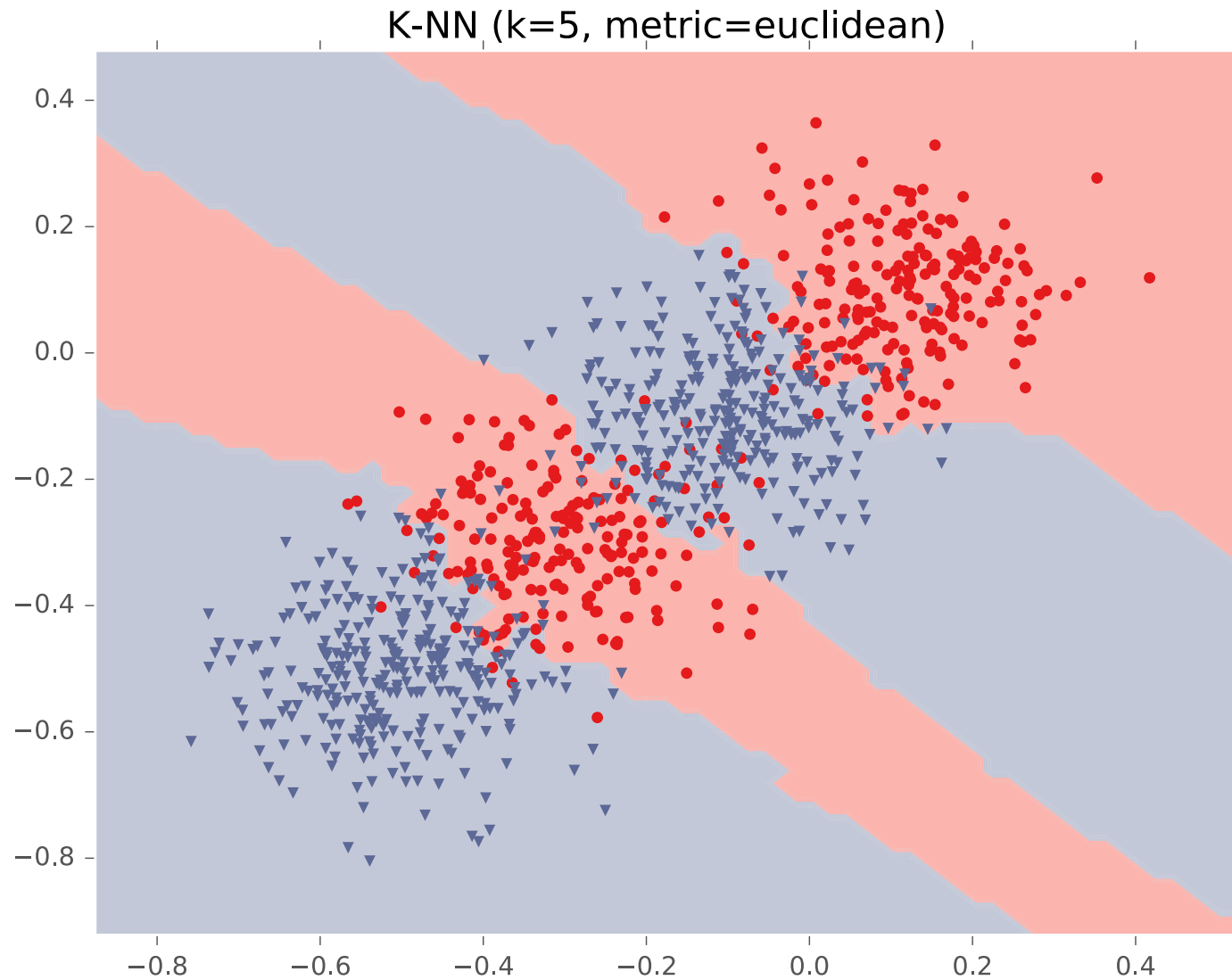
Example #3: Four Gaussians



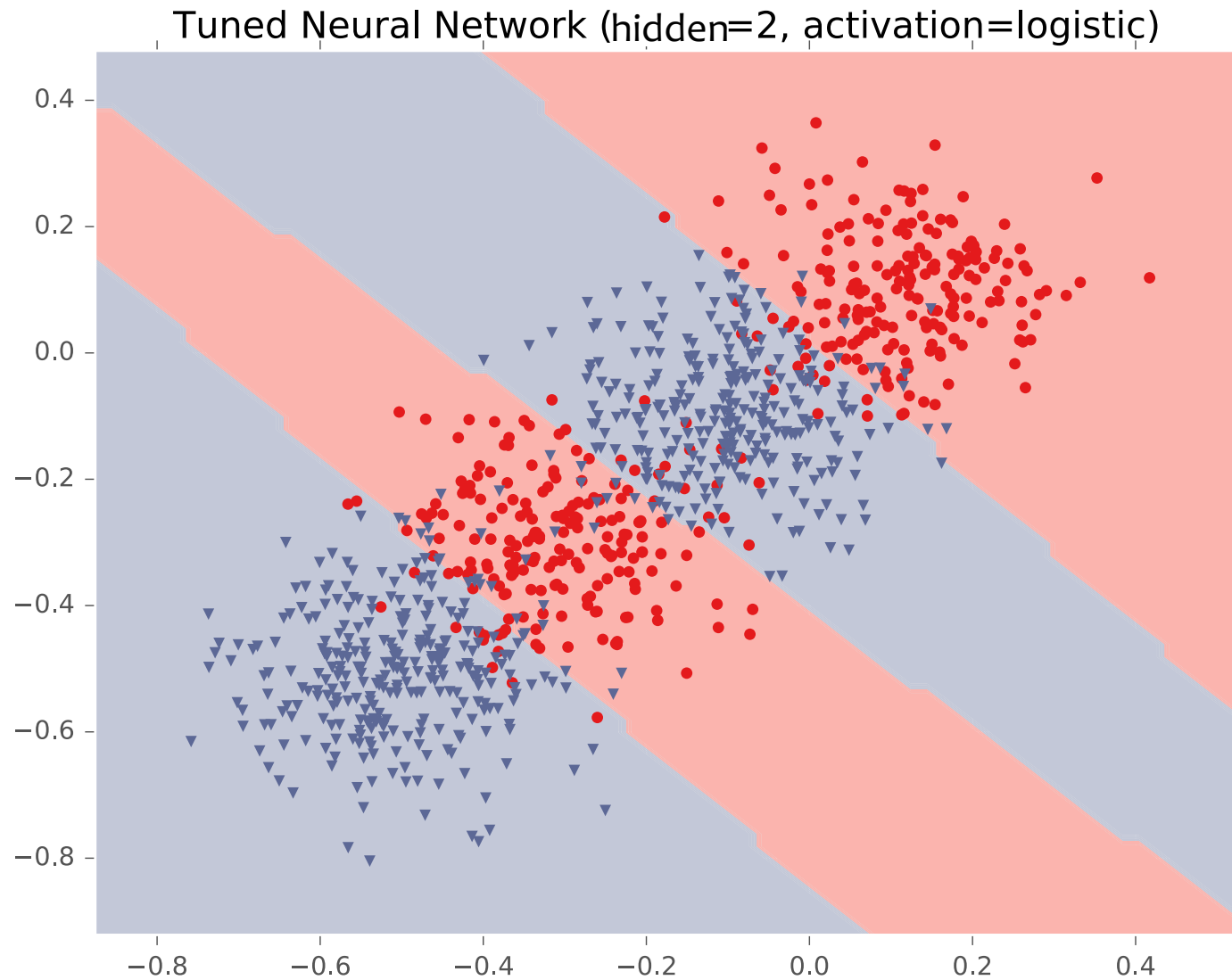
Example #3: Four Gaussians



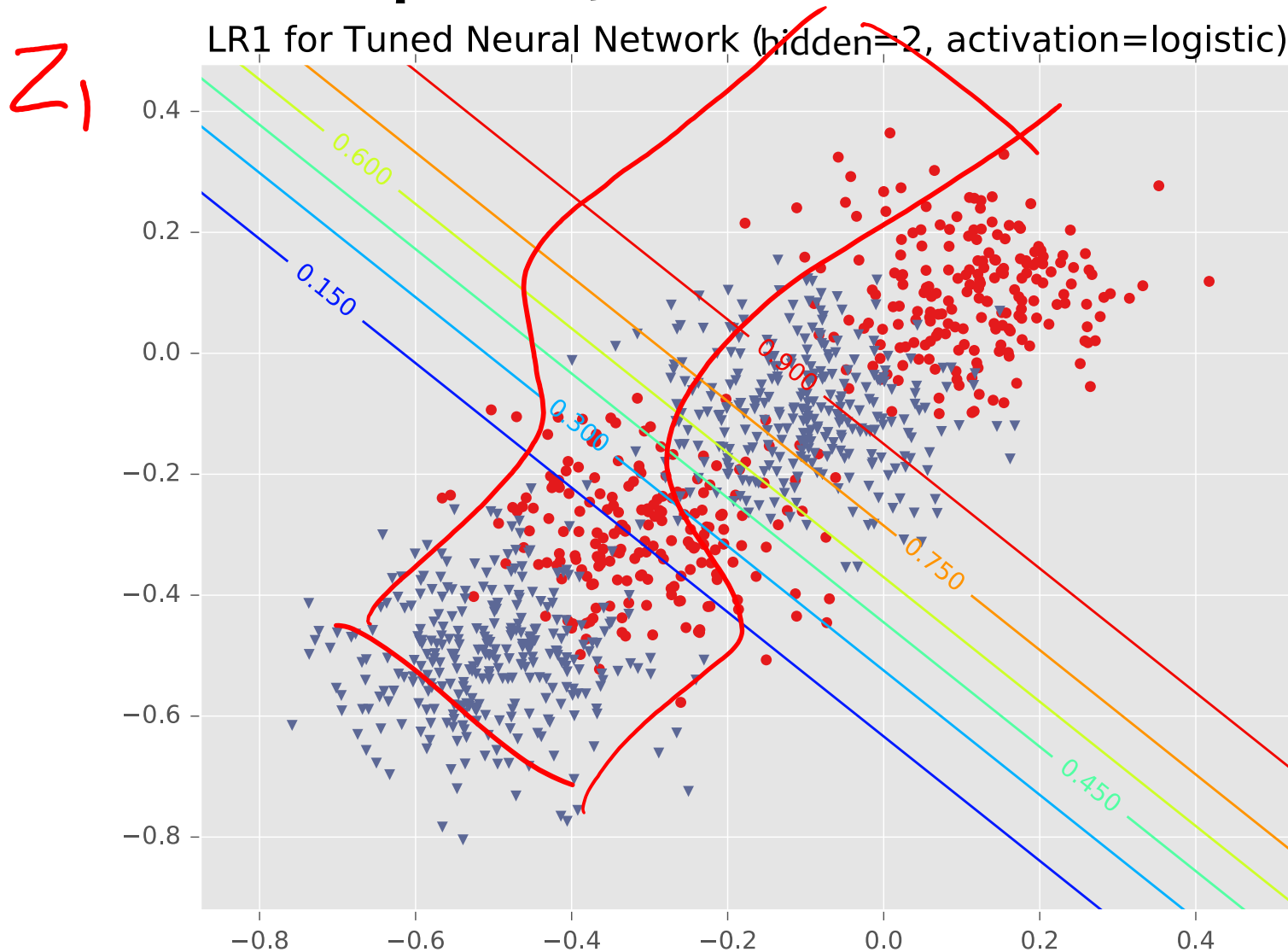
Example #3: Four Gaussians



Example #3: Four Gaussians

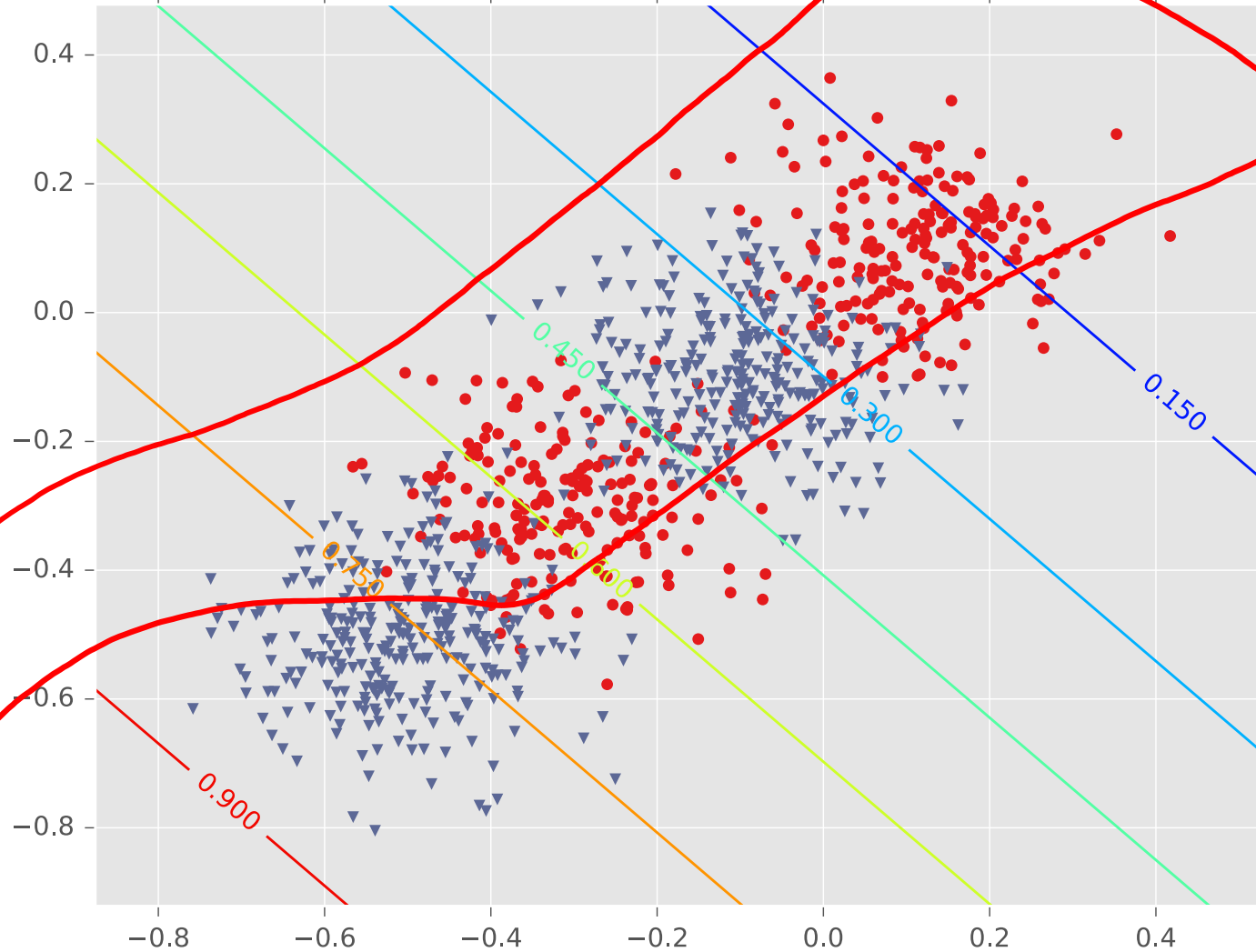


Example #3: Four Gaussians

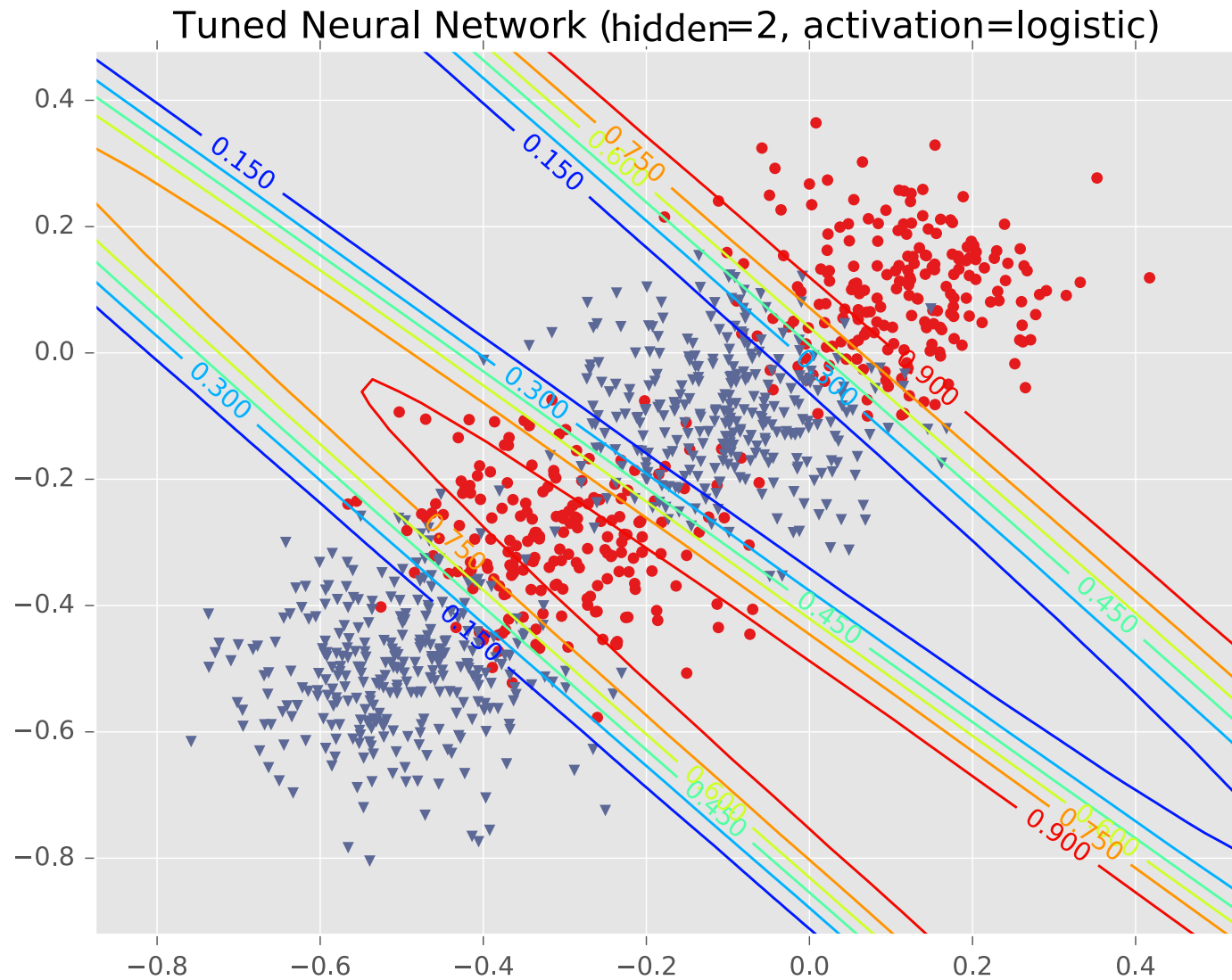


Example #3: Four Gaussians

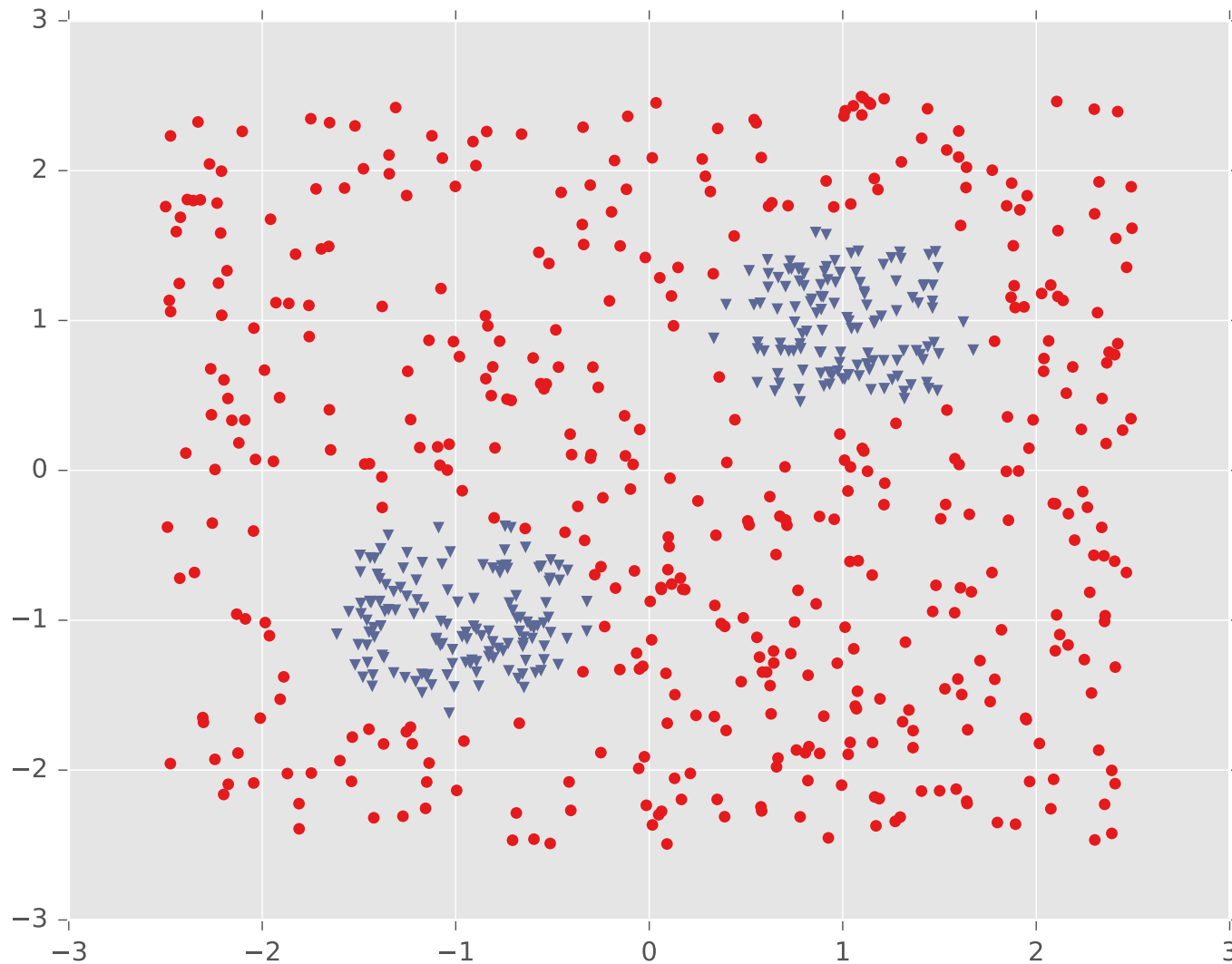
LR2 for Tuned Neural Network (hidden=2, activation=logistic)



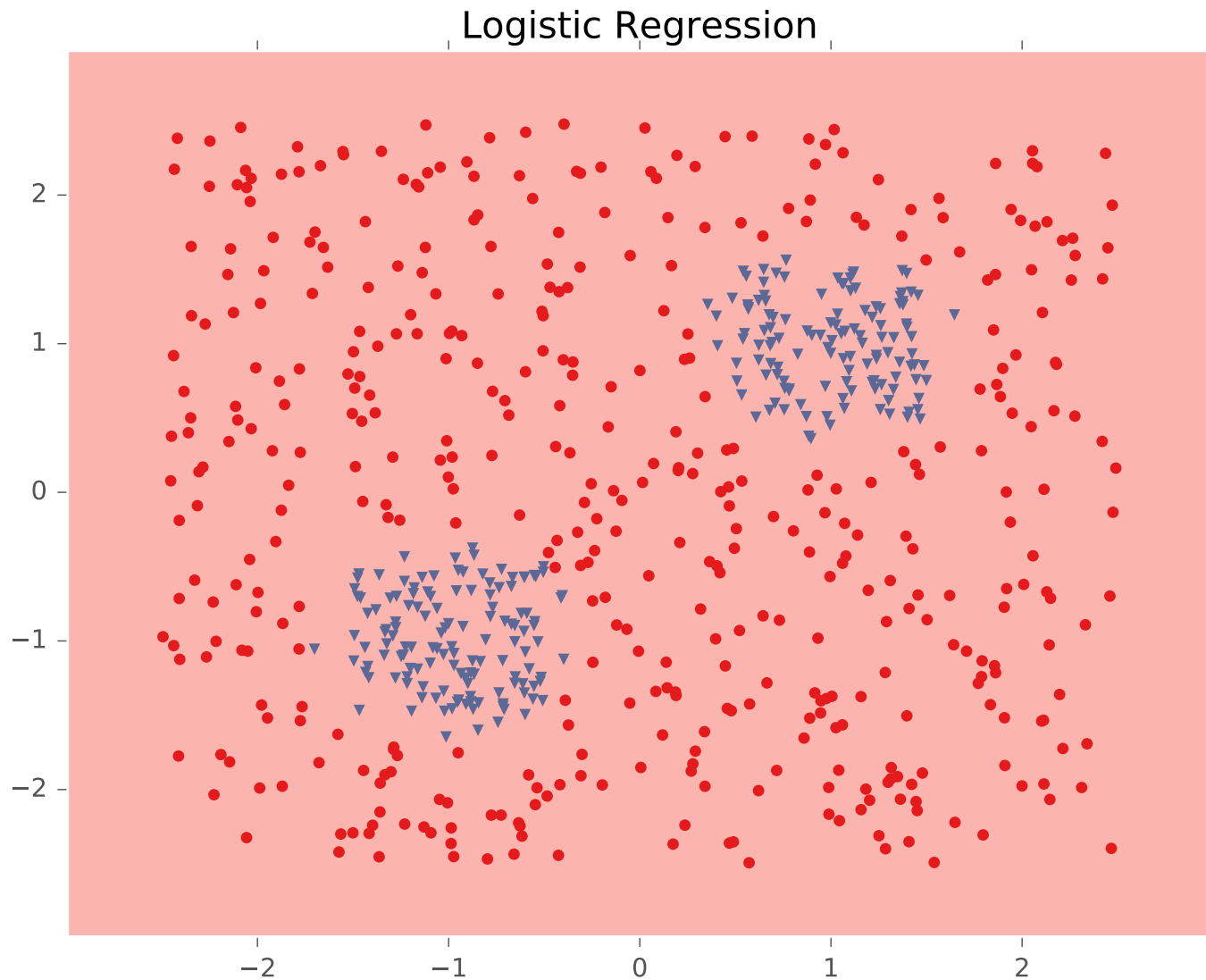
Example #3: Four Gaussians



Example #4: Two Pockets

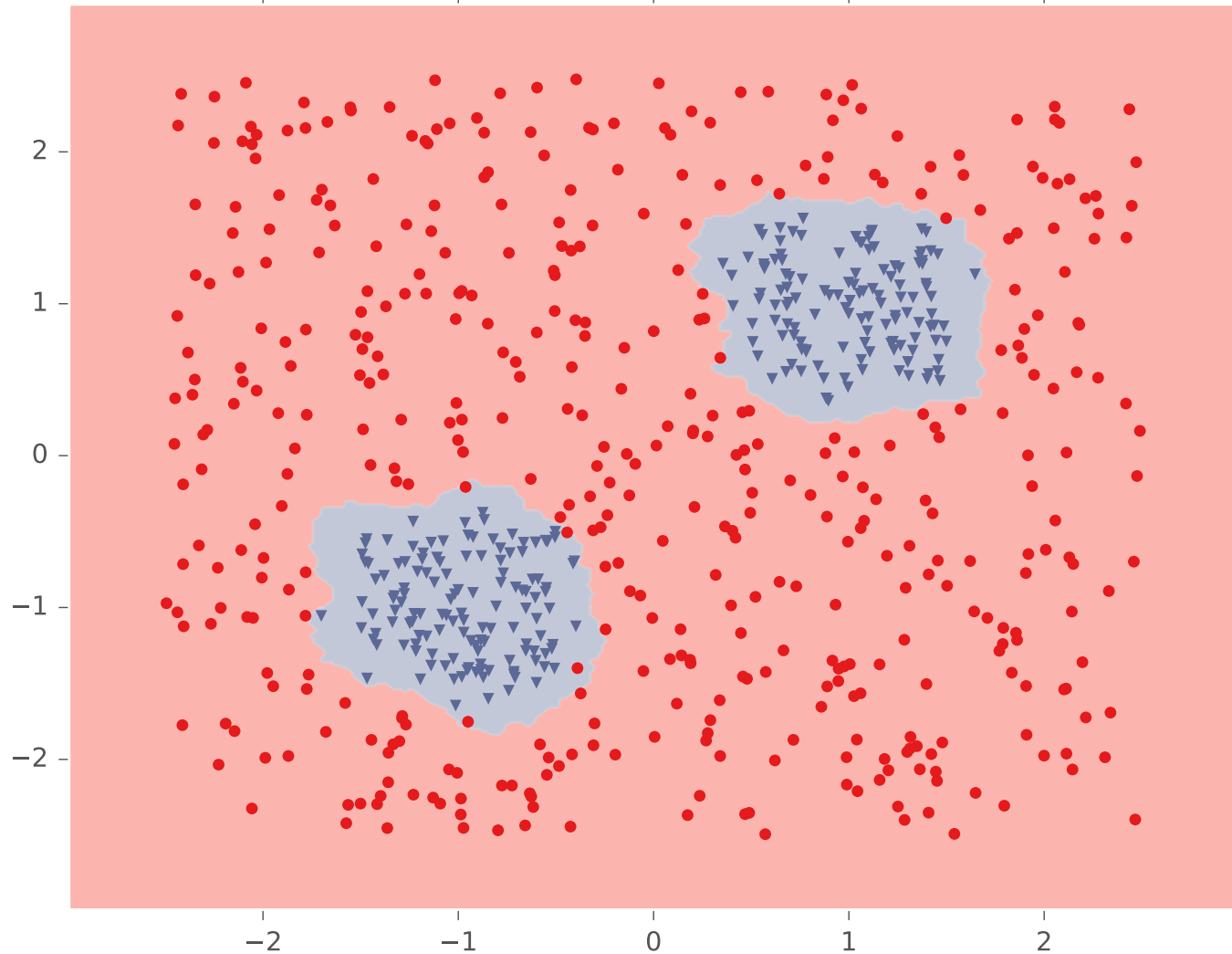


Example #4: Two Pockets



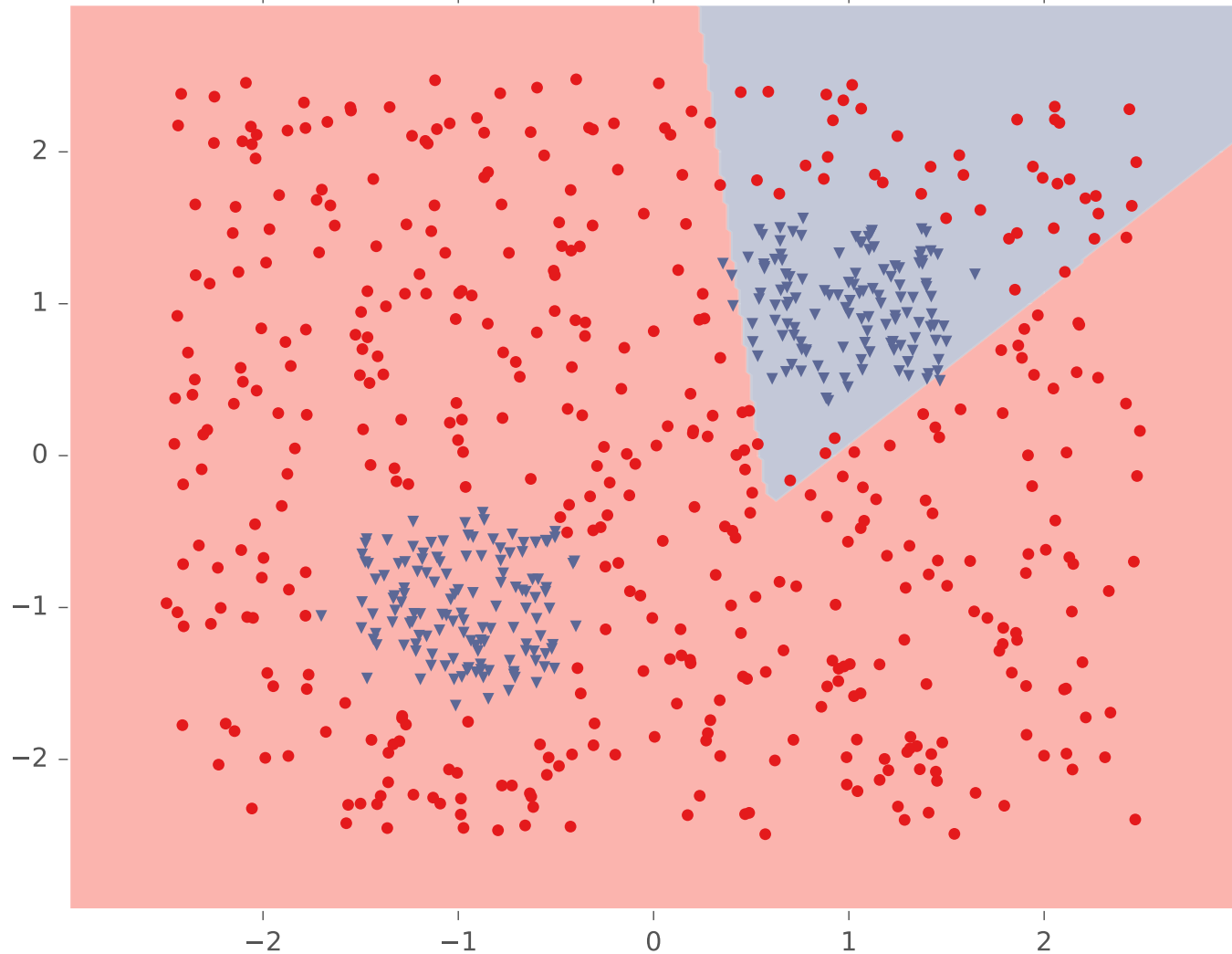
Example #4: Two Pockets

K-NN (k=5, metric=euclidean)



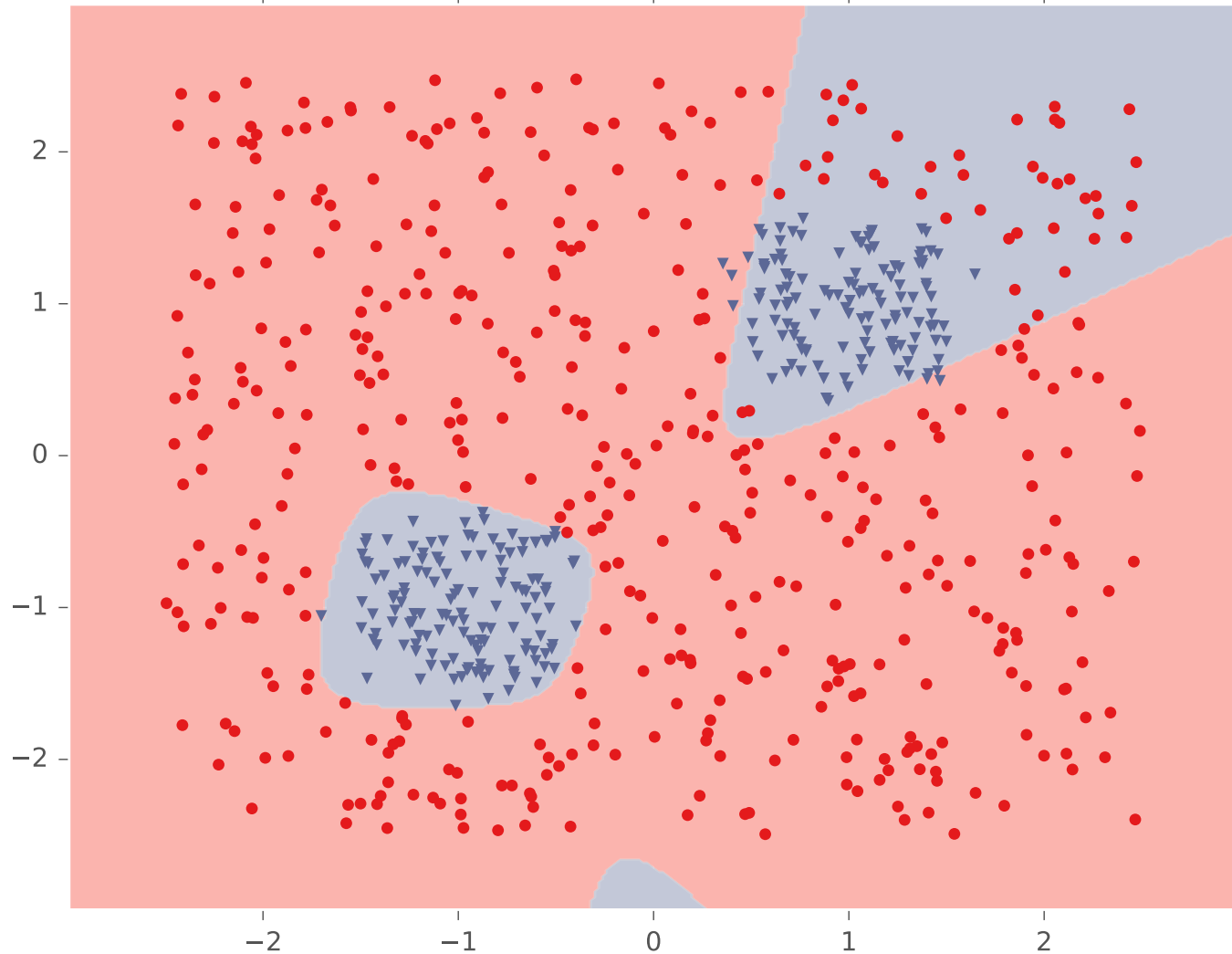
Example #4: Two Pockets

Tuned Neural Network (hidden=2, activation=logistic)



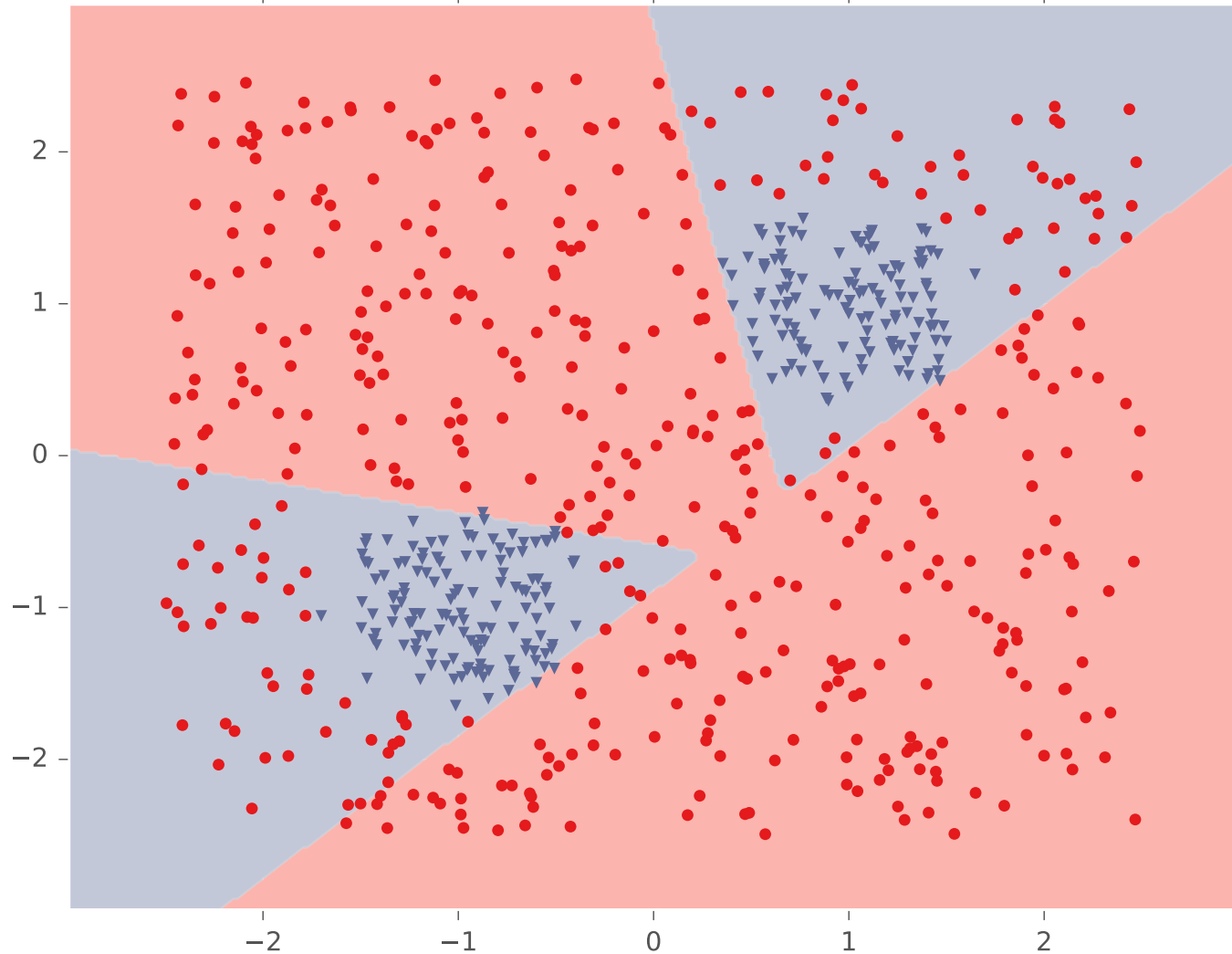
Example #4: Two Pockets

Tuned Neural Network (hidden=3, activation=logistic)



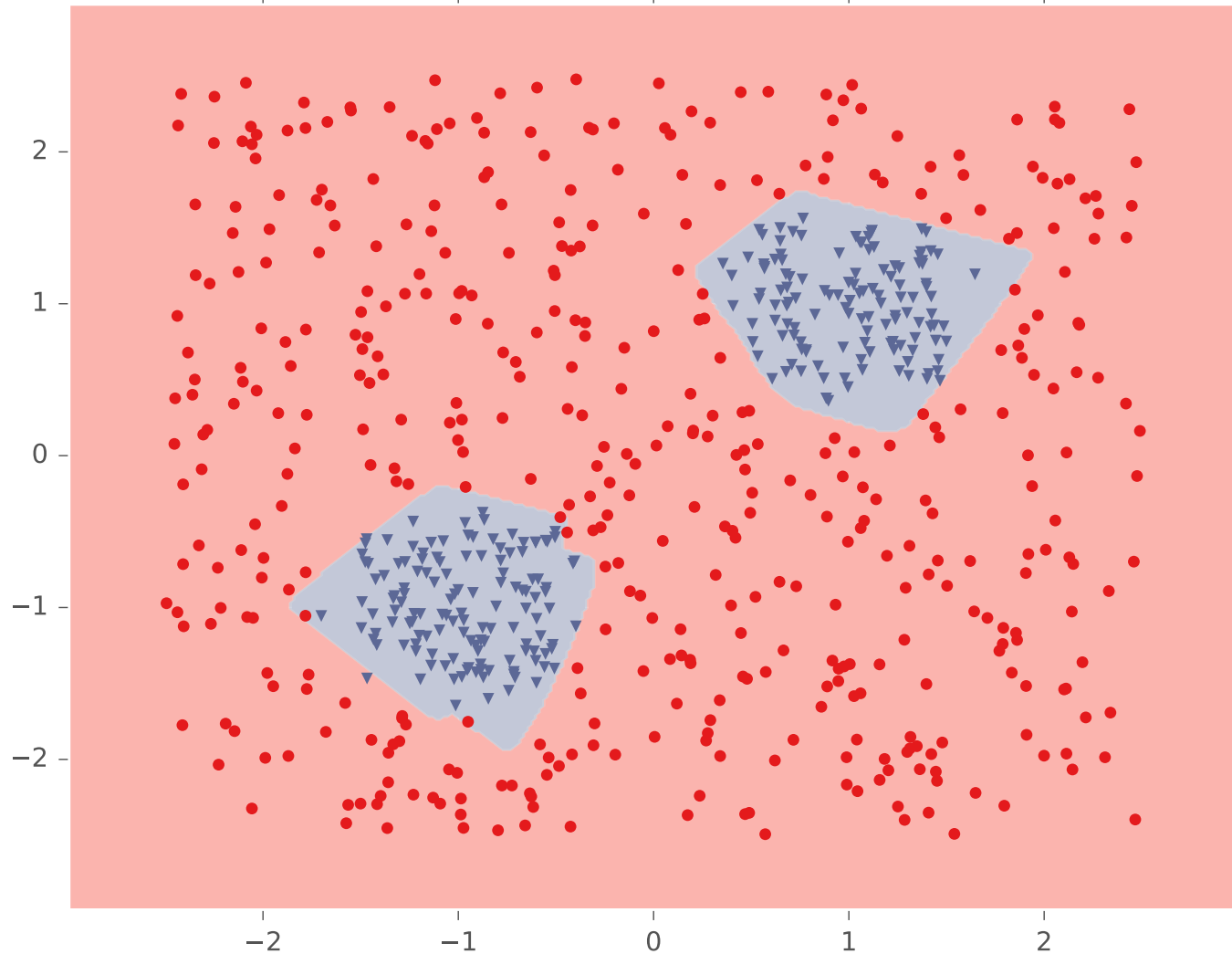
Example #4: Two Pockets

Tuned Neural Network (hidden=4, activation=logistic)



Example #4: Two Pockets

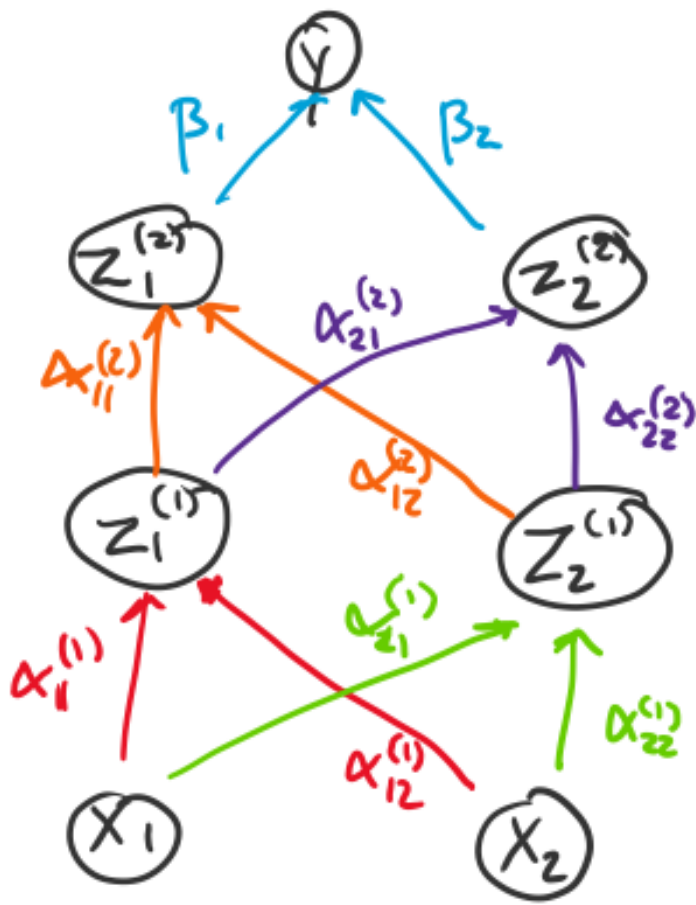
Tuned Neural Network (hidden=10, activation=logistic)



BUILDING DEEPER NETWORKS

Neural Net w/2 Hidden Layers

Ex#2: NN w/2 Hidden Layers and 2 units each



$$Z_1^{(1)} = \sigma(\alpha_{11}^{(1)} x_1 + \alpha_{12}^{(1)} x_2 + \alpha_{10}^{(1)})$$

$$Z_2^{(1)} = \sigma(\alpha_{21}^{(1)} x_1 + \alpha_{22}^{(1)} x_2 + \alpha_{20}^{(1)})$$

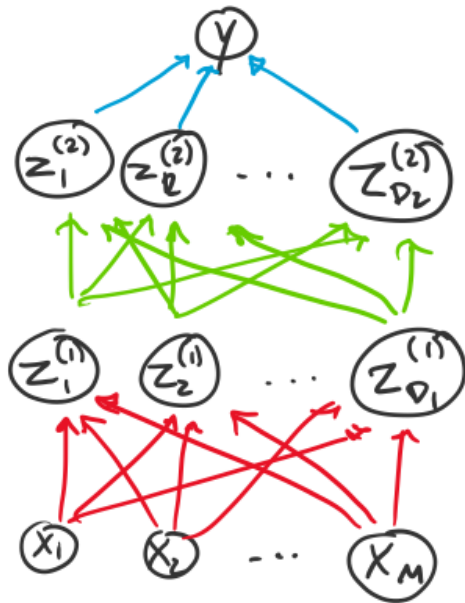
$$Z_1^{(2)} = \sigma(\alpha_{11}^{(2)} Z_1^{(1)} + \alpha_{12}^{(2)} Z_2^{(1)} + \alpha_{10}^{(2)})$$

$$Z_2^{(2)} = \sigma(\dots)$$

$$Y = \sigma(\dots)$$

Neural Net w/2 Hidden Layers

Ex#3: Arbitrary Feed-Forward NN



parameters

$$\beta \in \mathbb{R}^{D_2}$$

$$\beta_0 \in \mathbb{R}$$

$$\alpha^{(2)} \in \mathbb{R}^{D_1 \times D_2}$$

$$\vec{b}^{(2)} \in \mathbb{R}^{D_2}$$

$$\alpha^{(1)} \in \mathbb{R}^{M \times D_1}$$

$$\vec{b}^{(1)} \in \mathbb{R}^{D_1}$$

$$y = \sigma(\vec{\beta}^T \vec{z}^{(2)} + \beta_0)$$

$$\vec{z}^{(2)} = \sigma(\underbrace{(\alpha^{(2)})^T \vec{z}^{(1)} + \vec{b}^{(2)}}_{D_2 \times 1})$$

$$\vec{z}^{(1)} = \sigma(\underbrace{(\alpha^{(1)})^T \vec{x} + \vec{b}^{(1)}}_{\substack{D_1 \times M \quad M \times 1 \quad D_1 \times 1 \\ D_1 \times 1}})$$

we apply σ elementwise to each entry in the vector $\vec{u} \in \mathbb{R}^M$

$$\vec{z} = \sigma(\vec{u}) \in \mathbb{R}^M$$

$$z_j = \sigma(u_j)$$

Fold in the intercept terms?

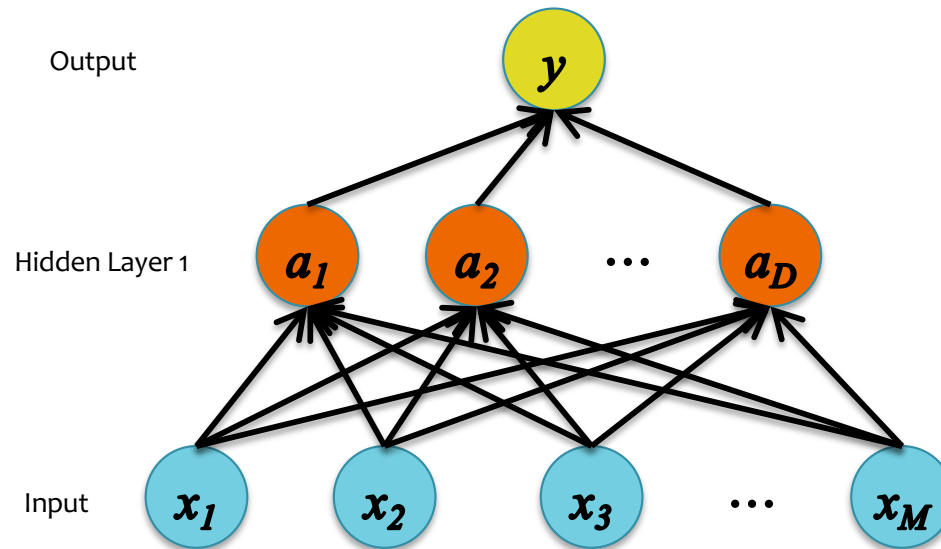
Assume $x_1 = 1, z_1^{(1)} = 1, z_1^{(2)} = 1$.

Drop $\vec{b}^{(1)}, \vec{b}^{(2)}, \beta_0$

Caution: tricky to implement

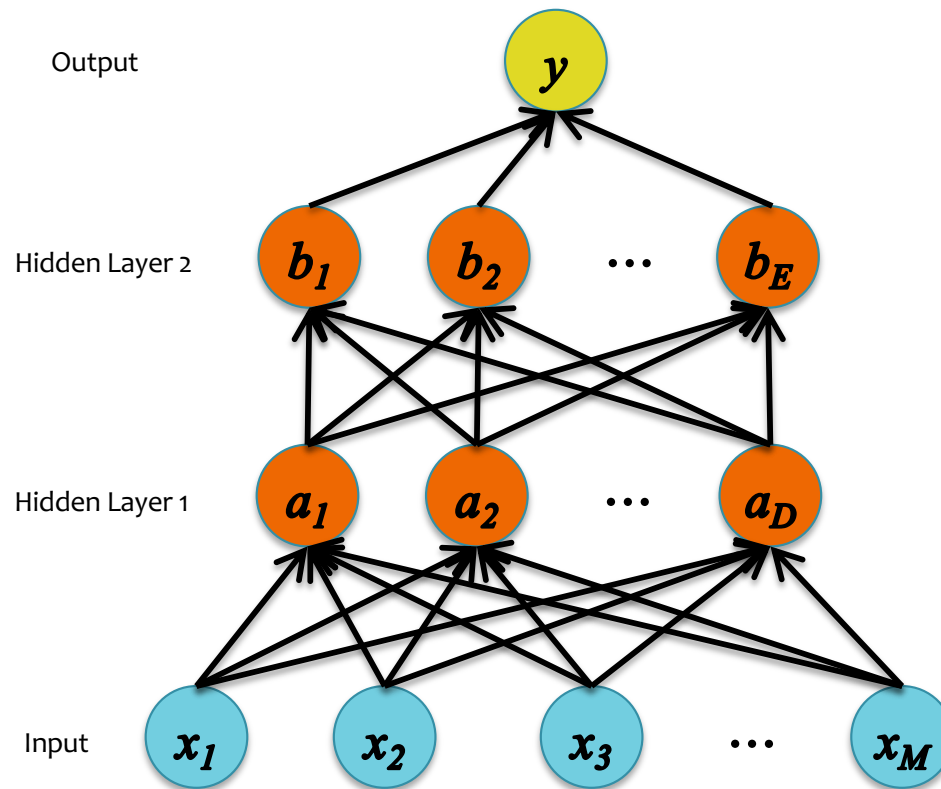
Deeper Networks

Q: How many layers should we use?



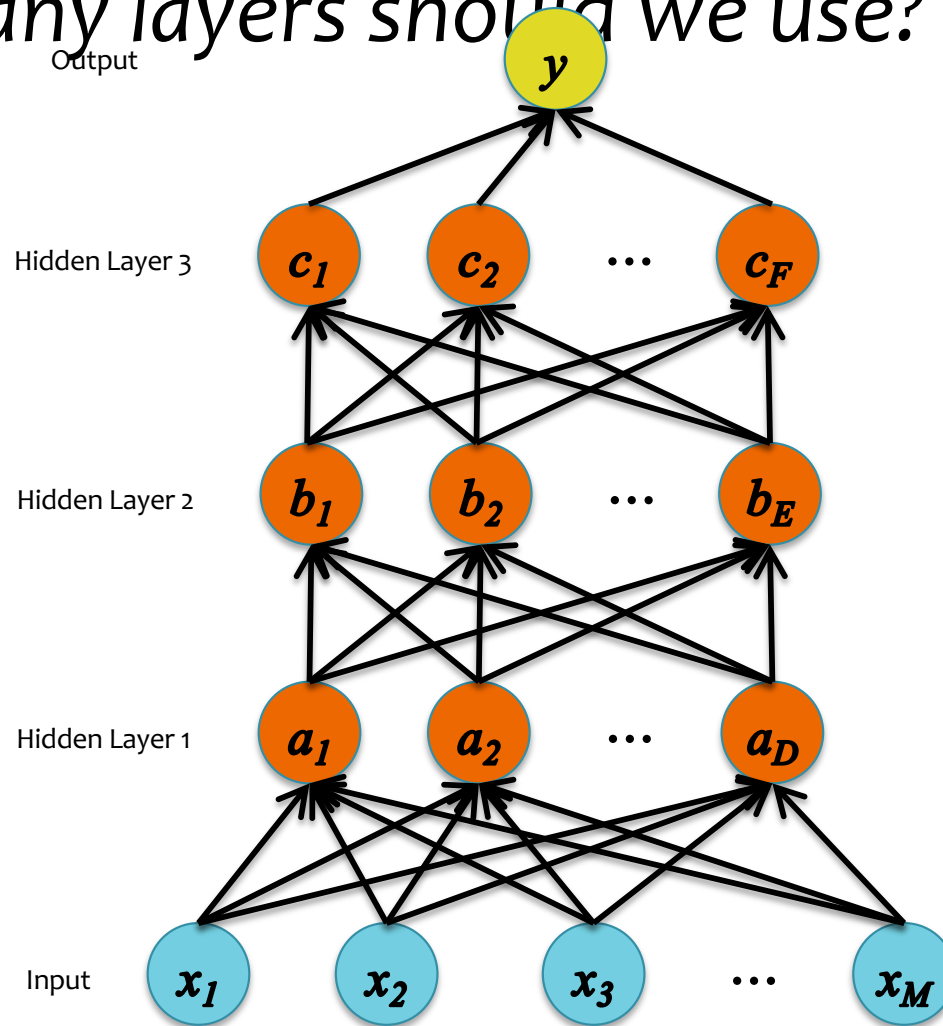
Deeper Networks

Q: How many layers should we use?



Deeper Networks

Q: *How many layers should we use?*



Deeper Networks

Q: How many layers should we use?

- **Theoretical answer:**

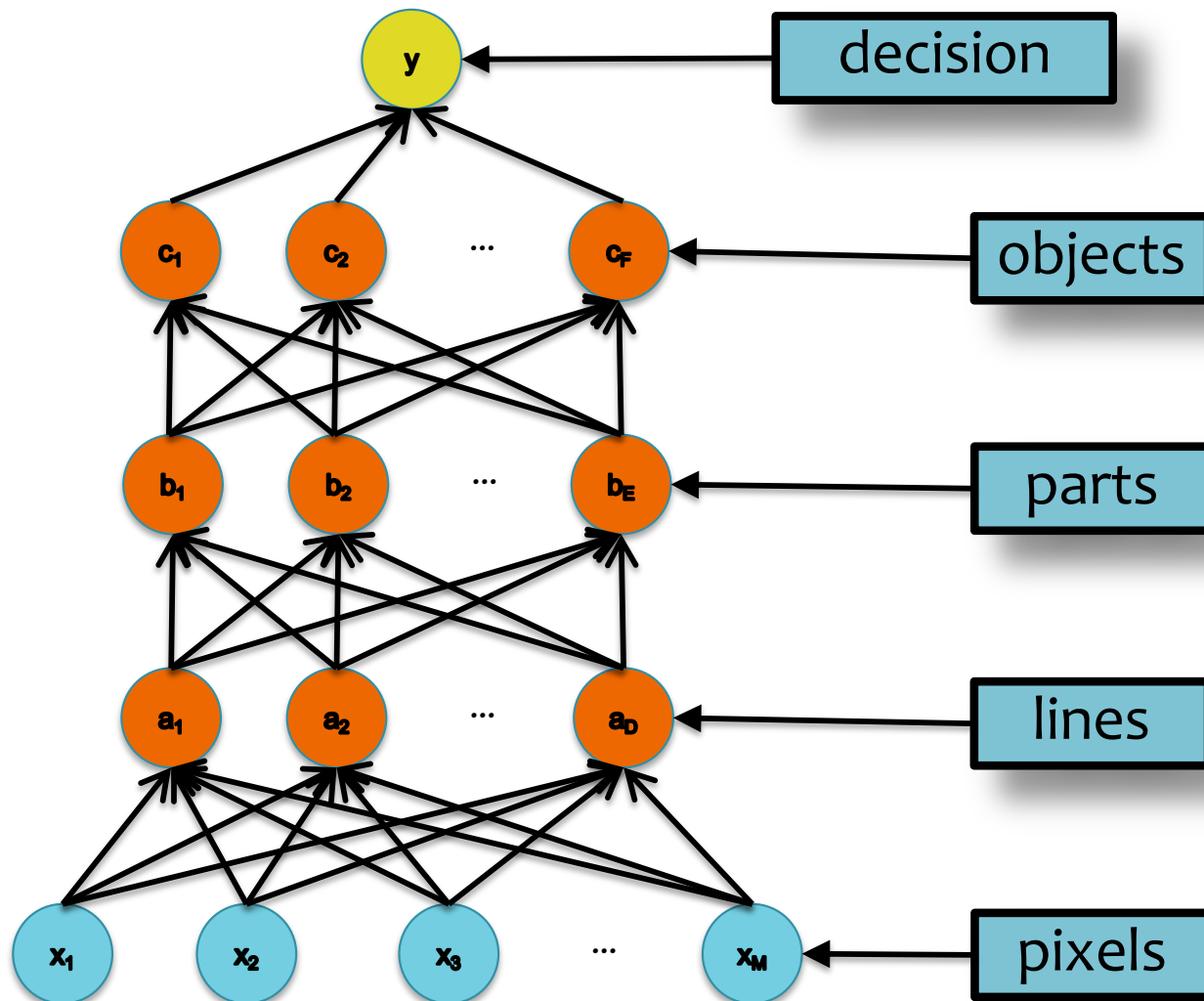
- A neural network with 1 hidden layer is a **universal function approximator**
- Cybenko (1989): For any continuous function $g(\mathbf{x})$, there exists a 1-hidden-layer neural net $h_{\theta}(\mathbf{x})$ s.t. $|h_{\theta}(\mathbf{x}) - g(\mathbf{x})| < \epsilon$ for all $\mathbf{x}$, assuming sigmoid activation functions

- **Empirical answer:**

- Before 2006: “Deep networks (e.g. 3 or more hidden layers) are too hard to train”
- After 2006: “Deep networks are easier to train than shallow networks (e.g. 2 or fewer layers) for many problems”

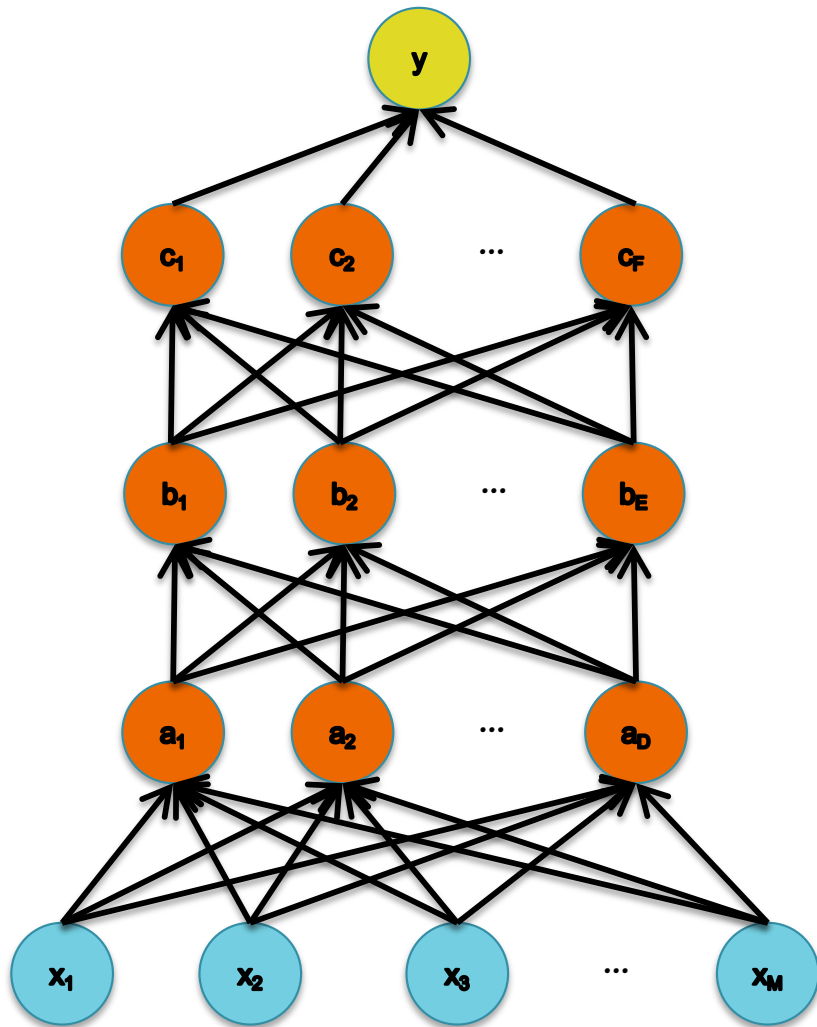
Big caveat: You need to know and use the right tricks.

Feature Learning



- **Traditional feature engineering**: build up levels of abstraction by hand
- **Deep networks** (e.g. convolution networks): learn the increasingly higher levels of abstraction from data
 - each layer is a learned feature representation
 - sophistication increases in higher layers

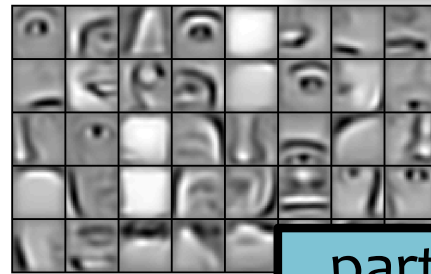
Feature Learning



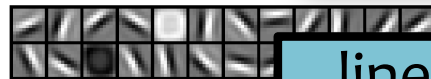
CBDN on Faces



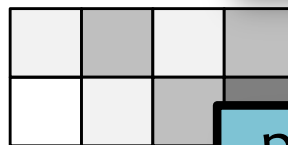
objects



parts



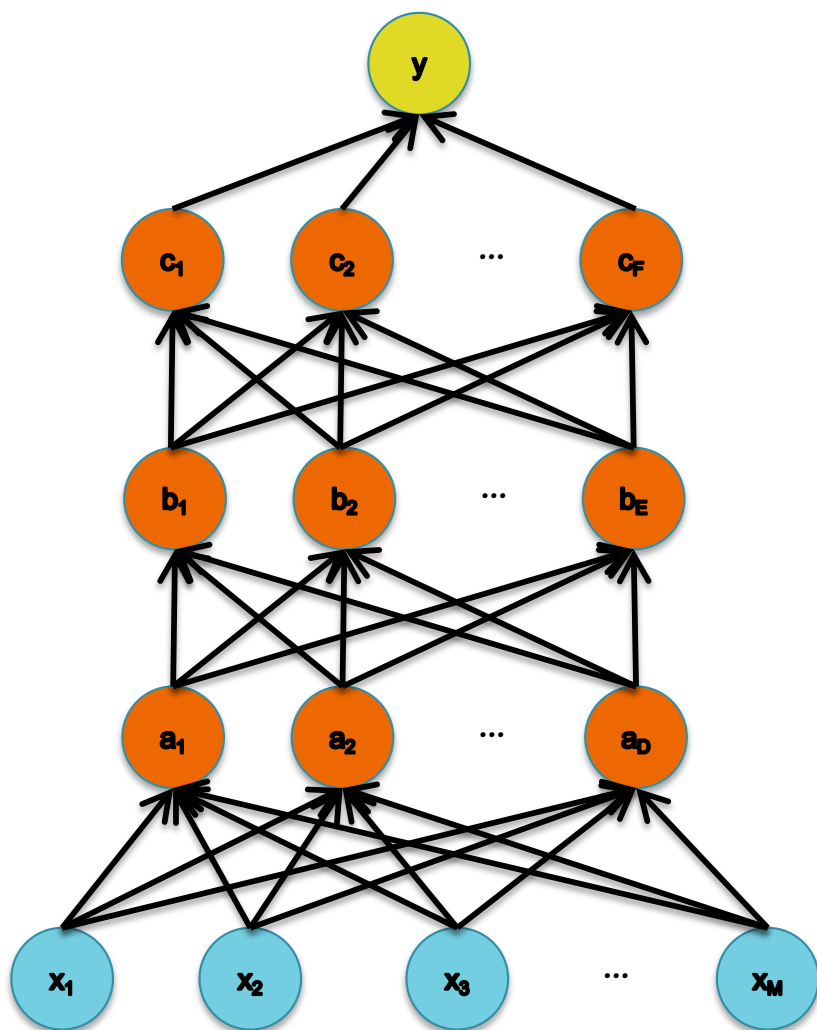
lines



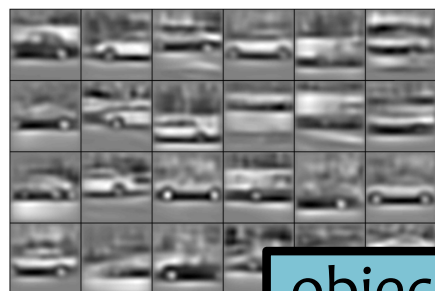
pixels

- **Traditional feature engineering:** build up levels of abstraction by hand
- **Deep networks** (e.g. convolution networks): learn the increasingly higher levels of abstraction from data
 - each layer is a learned feature representation
 - sophistication increases in higher layers

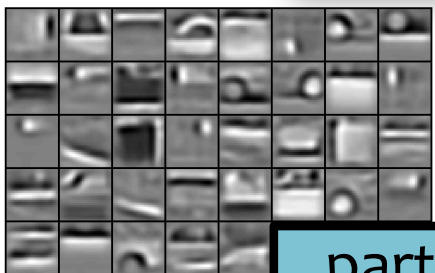
Feature Learning



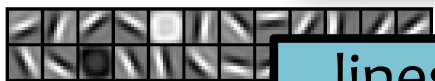
CBDN on Cars



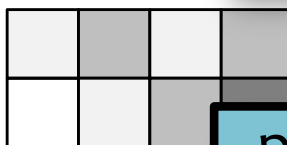
objects



parts



lines



pixels

- **Traditional feature engineering:** build up levels of abstraction by hand
- **Deep networks** (e.g. convolution networks): learn the increasingly higher levels of abstraction from data
 - each layer is a learned feature representation
 - sophistication increases in higher layers

ARCHITECTURES

Neural Network Architectures

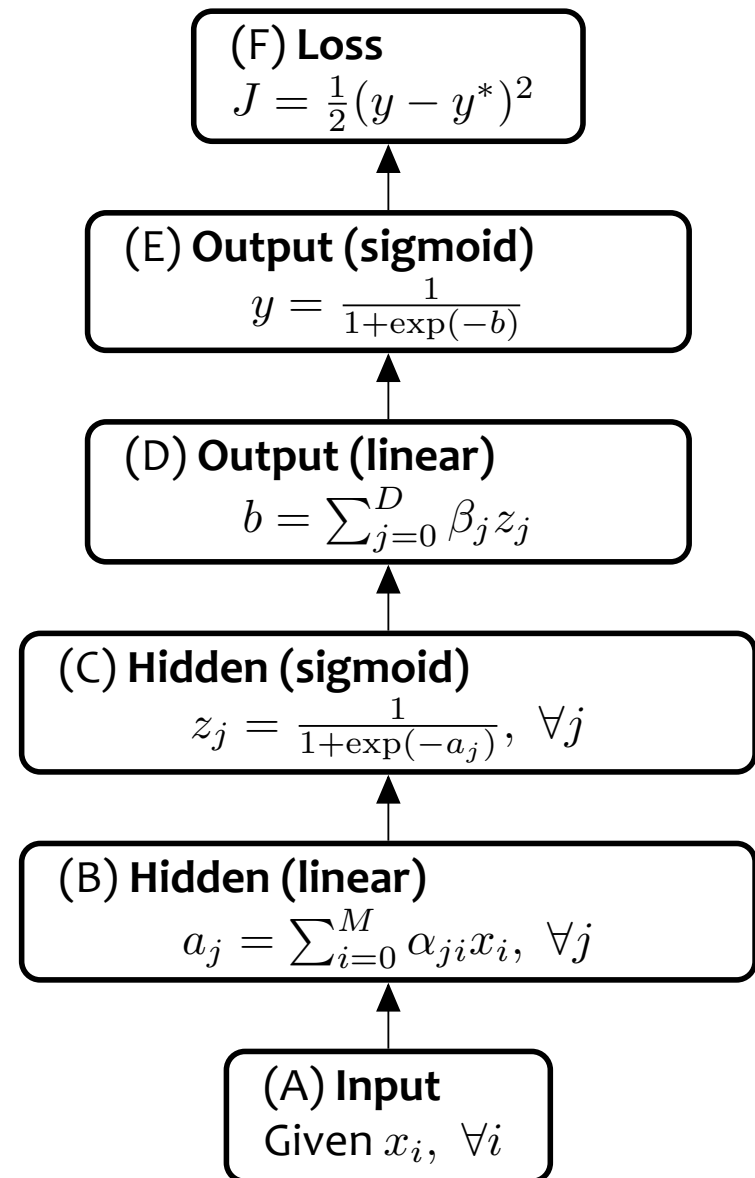
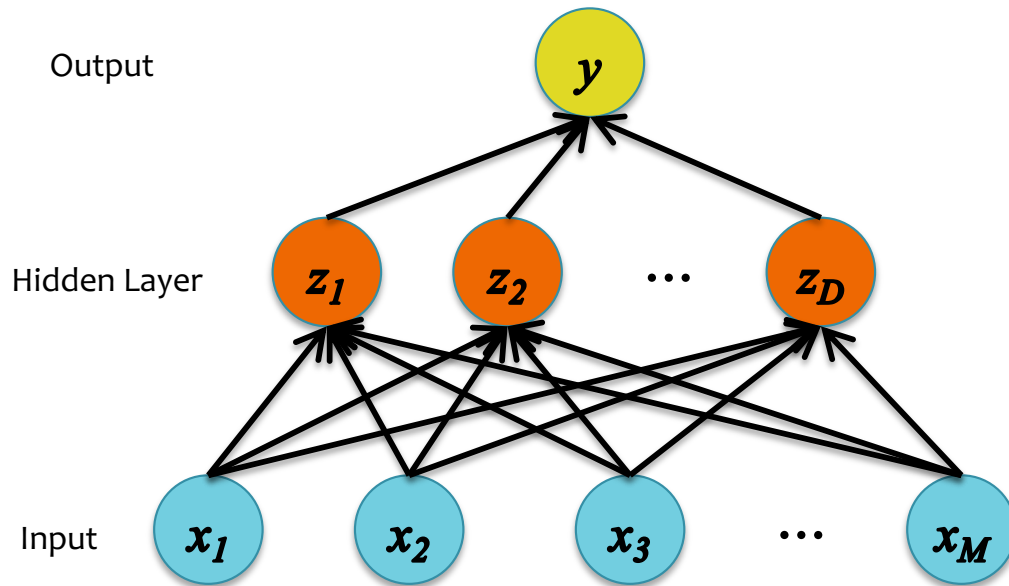
Even for a basic Neural Network, there are many design decisions to make:

1. # of hidden layers (depth)
2. # of units per hidden layer (width)
3. Type of activation function (nonlinearity)
4. Form of objective function
5. How to initialize the parameters

ACTIVATION FUNCTIONS

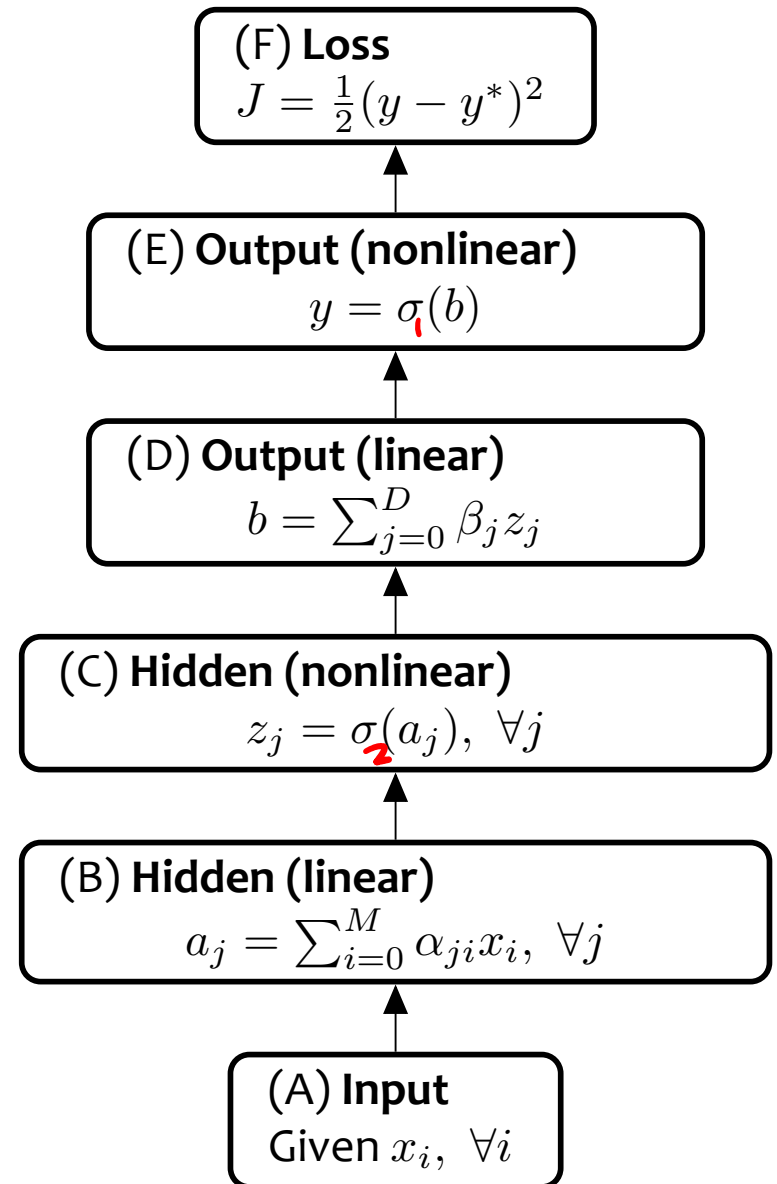
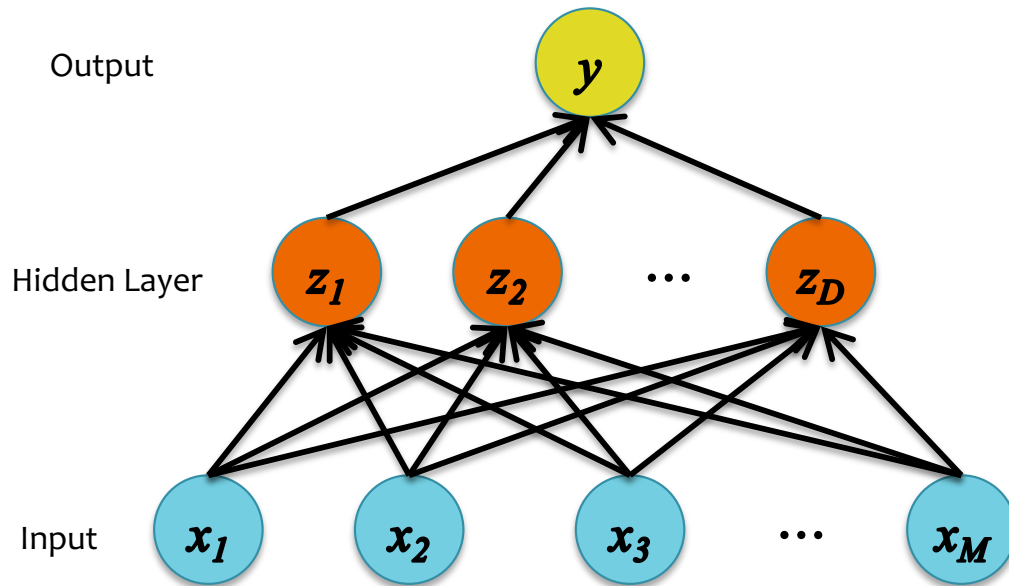
Activation Functions

Neural Network with sigmoid
activation functions



Activation Functions

Neural Network with arbitrary nonlinear activation functions

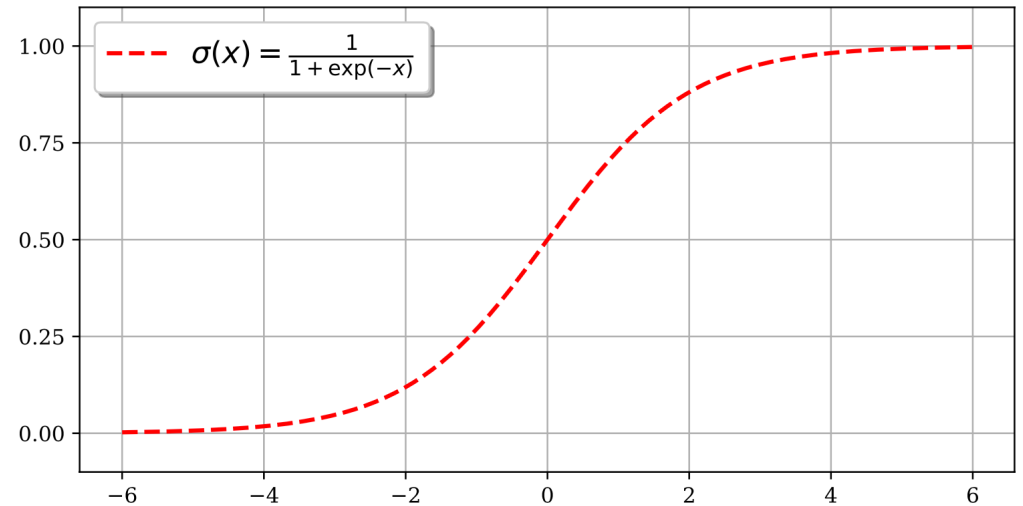


Activation Functions

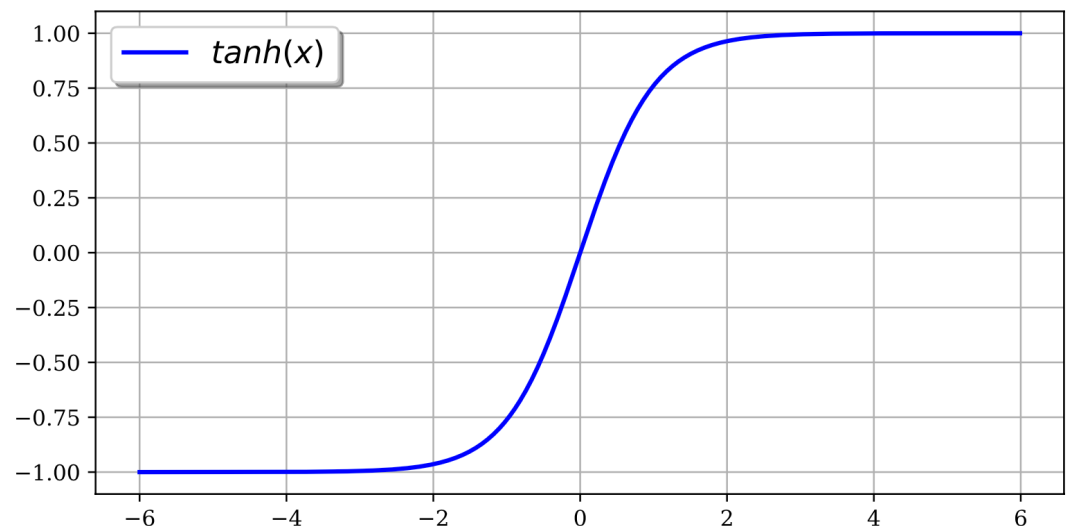
So far, we've assumed that the activation function (nonlinearity) is always the sigmoid function...

... but the sigmoid is not widely used in modern neural networks

Sigmoid (aka. logistic) function



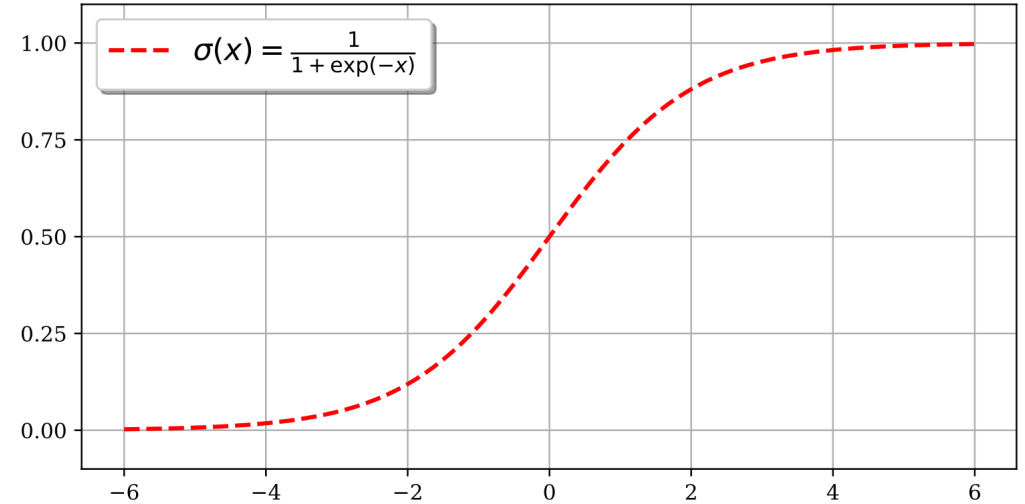
Hyperbolic tangent function



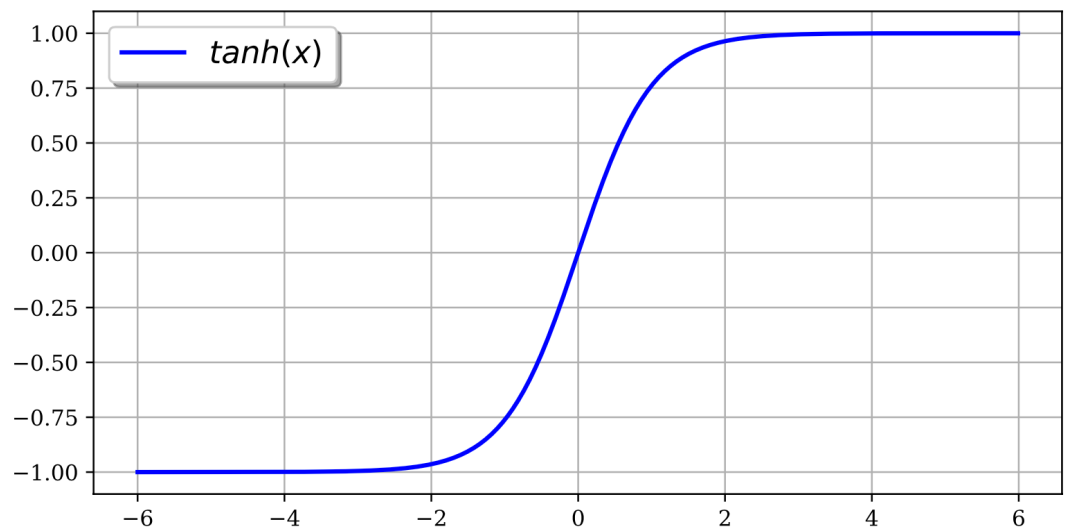
Activation Functions

- sigmoid, $\sigma(x)$
 - output in range (0,1)
 - good for probabilistic outputs
- hyperbolic tangent, $\tanh(x)$
 - similar shape to sigmoid, but output in range (-1,+1)

Sigmoid (aka. logistic) function

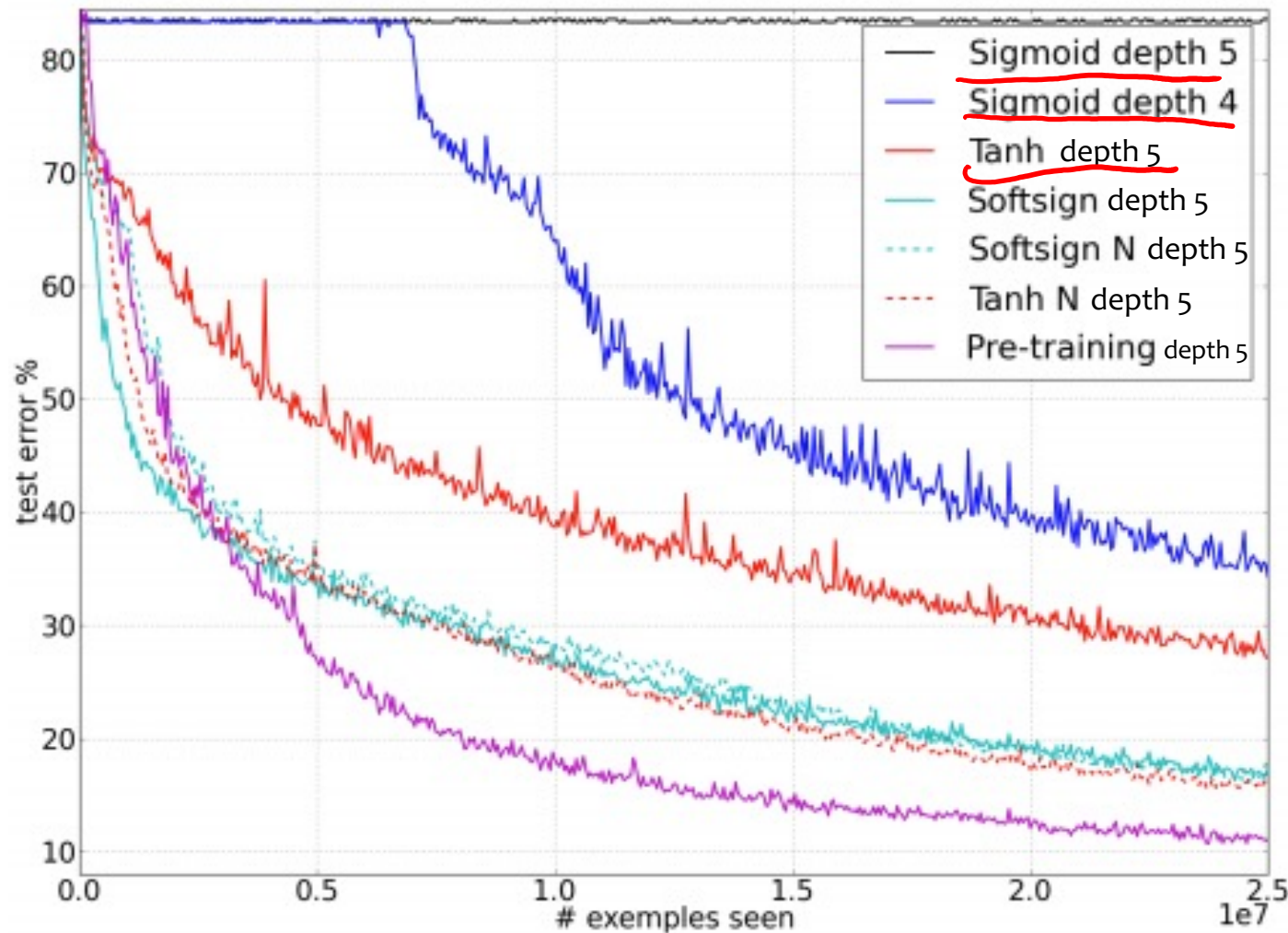


Hyperbolic tangent function



Understanding the difficulty of training deep feedforward neural networks

AI Stats 2010



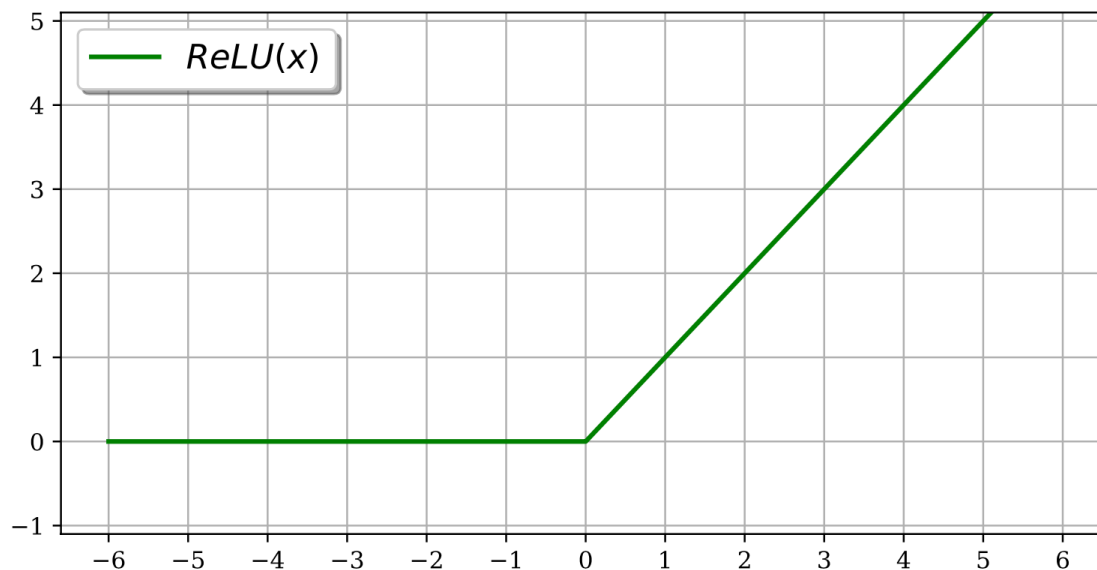
} sigmoid
vs.
tanh

Figure from Glorot & Bentio (2010)

Activation Functions

- Rectified Linear Unit (ReLU)
 - avoids the vanishing gradient problem
 - derivative is fast to compute

$$\text{ReLU}(x) = \max(0, x)$$

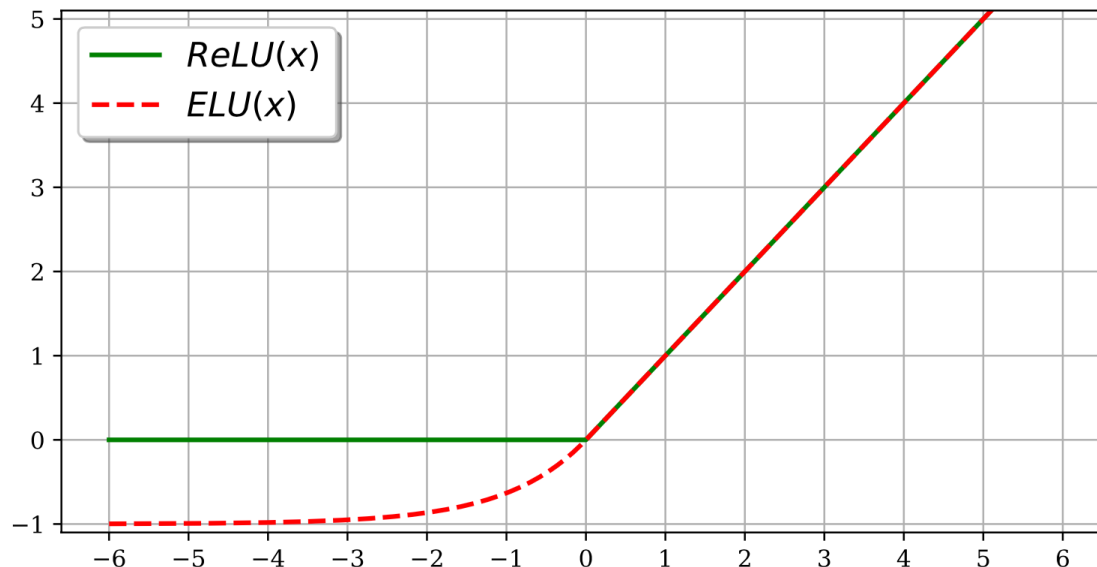


Activation Functions

- Rectified Linear Unit (ReLU)

- avoids the vanishing gradient problem
- derivative is fast to compute

$$\text{ReLU}(x) = \max(0, x)$$



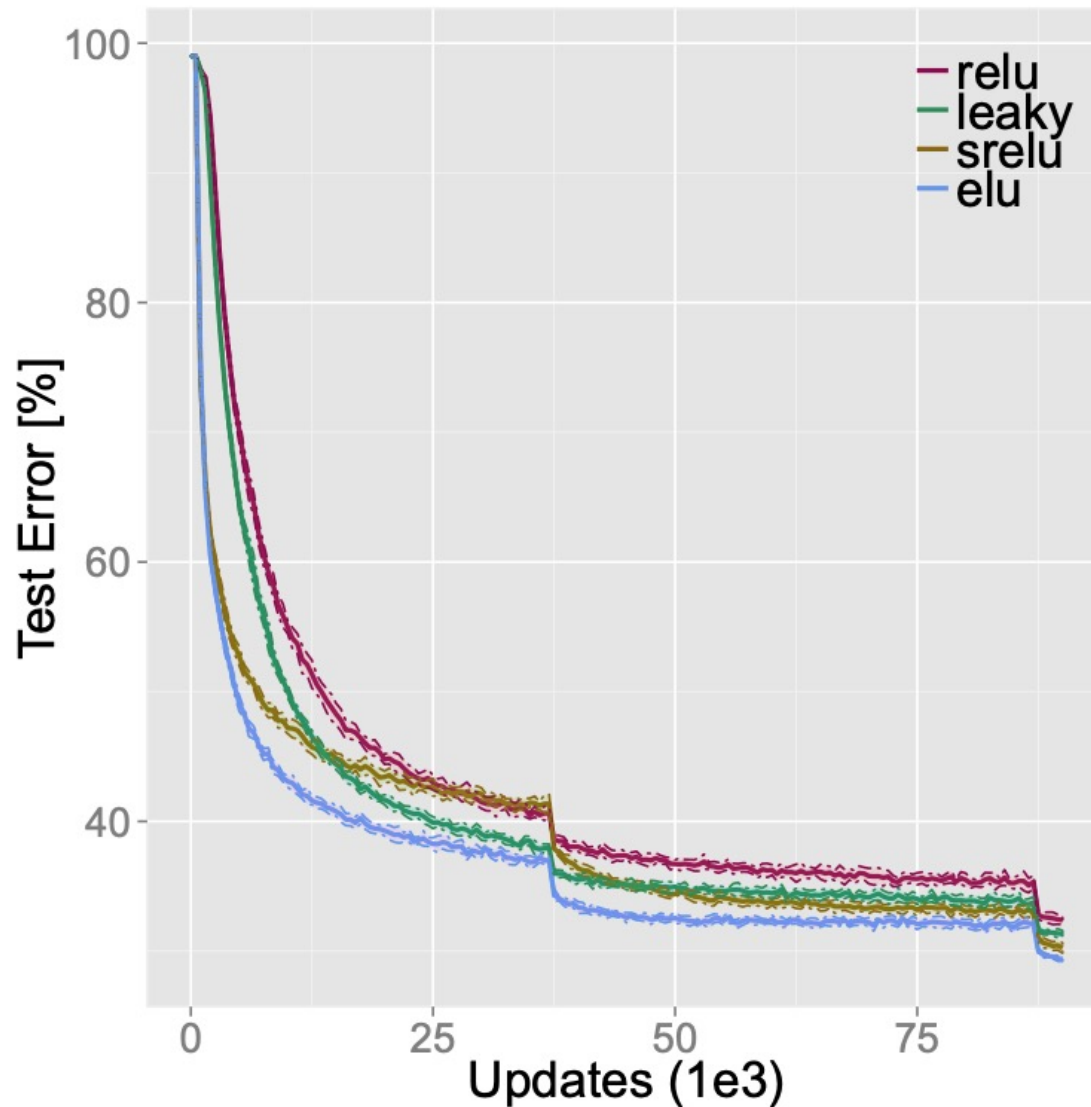
- Exponential Linear Unit (ELU)

- same as ReLU on positive inputs
- unlike ReLU, allows negative outputs and smoothly transitions for $x < 0$

$$\text{ELU}(x) = \begin{cases} x, & \text{if } x > 0 \\ \alpha(\exp(x) - 1), & \text{if } x \leq 0 \end{cases}$$

Activation Functions

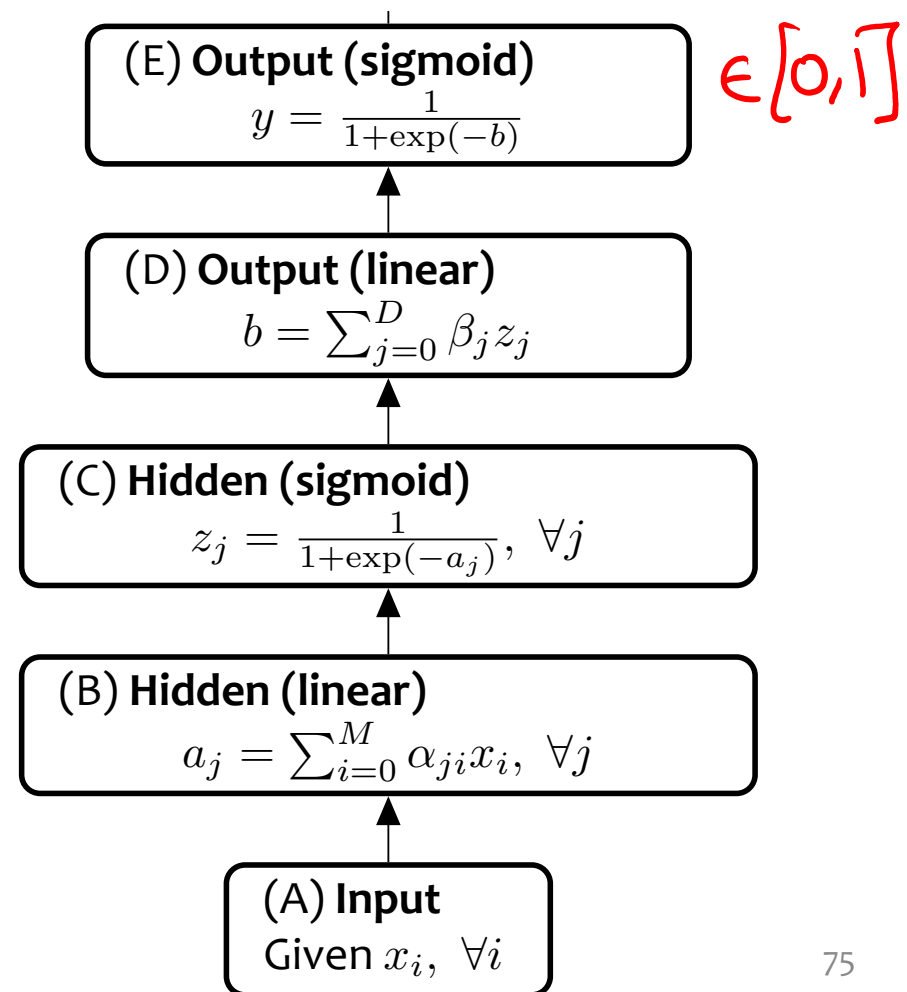
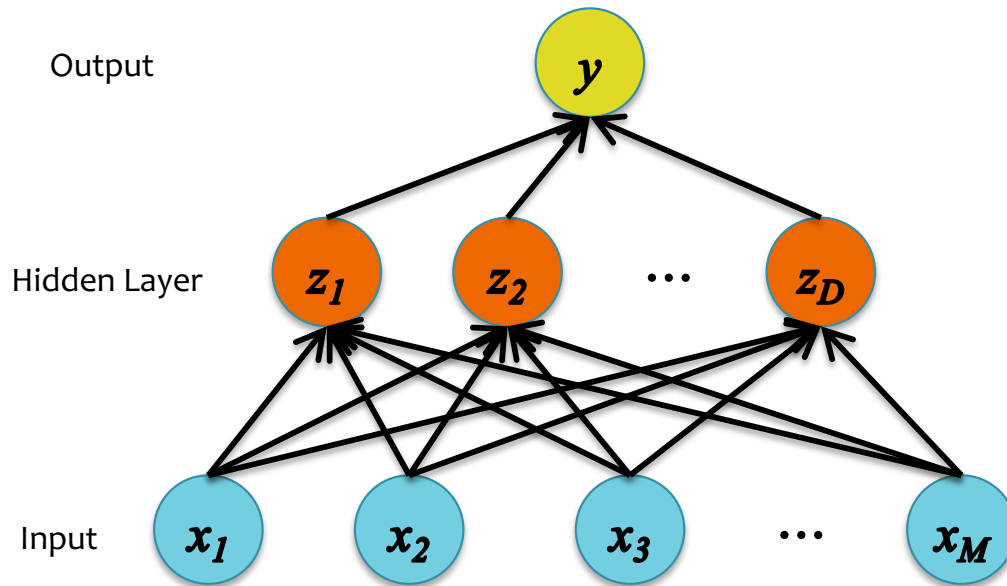
Image Classification Benchmark (CIFAR-10)



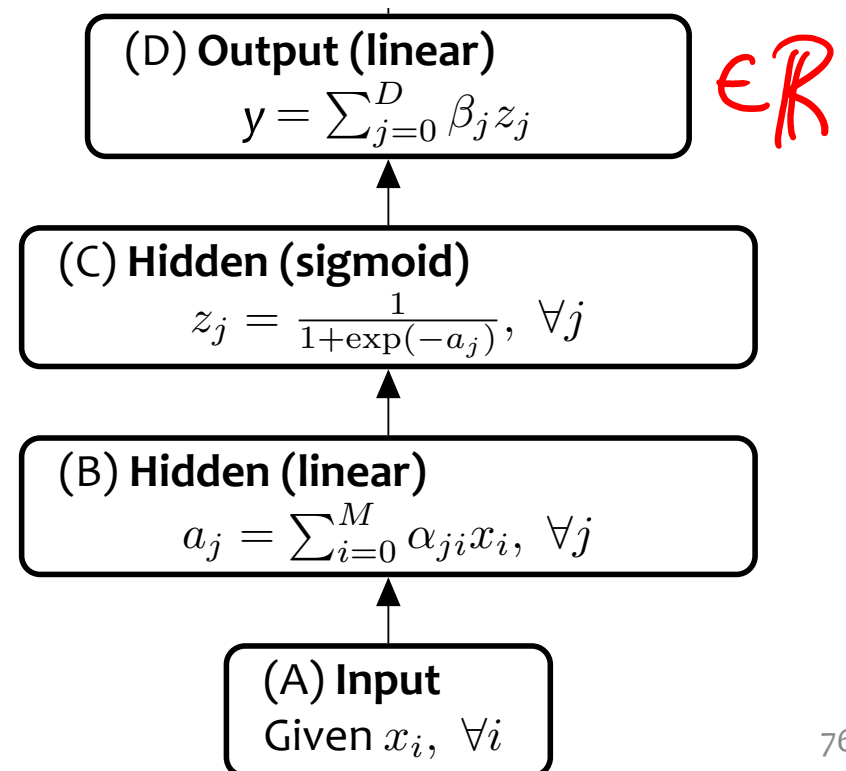
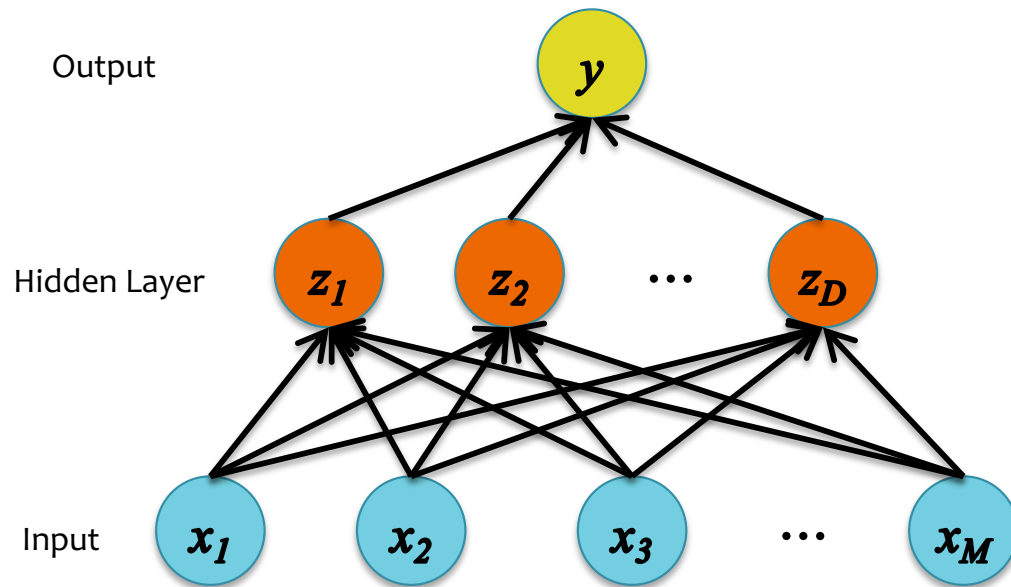
1. Training loss converges fastest with ELU
2. ELU(x) yields lower test error than ReLU(x) on CIFAR-10

LOSS FUNCTIONS & OUTPUT LAYERS

Neural Network for Classification



Neural Network for Regression



Objective Functions for NNs

1. Quadratic Loss:

- the same objective as Linear Regression
- i.e. mean squared error

$$J = \ell_Q(y, y^{(i)}) = \frac{1}{2}(y - y^{(i)})^2$$
$$\frac{dJ}{dy} = y - y^{(i)}$$

2. Binary Cross-Entropy:

- the same objective as Binary Logistic Regression
- i.e. negative log likelihood
- This requires our output y to be a probability in $[0,1]$

$$J = \ell_{CE}(y, y^{(i)}) = -(y^{(i)} \log(y) + (1 - y^{(i)}) \log(1 - y))$$
$$\frac{dJ}{dy} = - \left(y^{(i)} \frac{1}{y} + (1 - y^{(i)}) \frac{1}{y - 1} \right)$$

Objective Functions for NNs

Cross-entropy vs. Quadratic loss

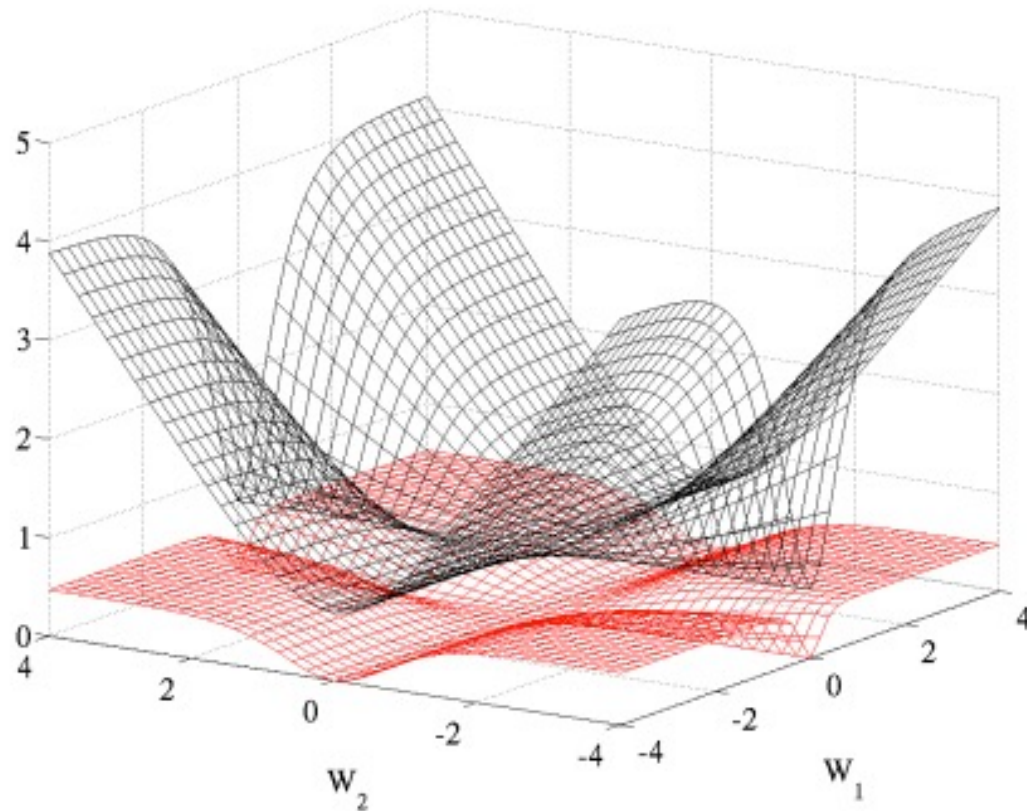
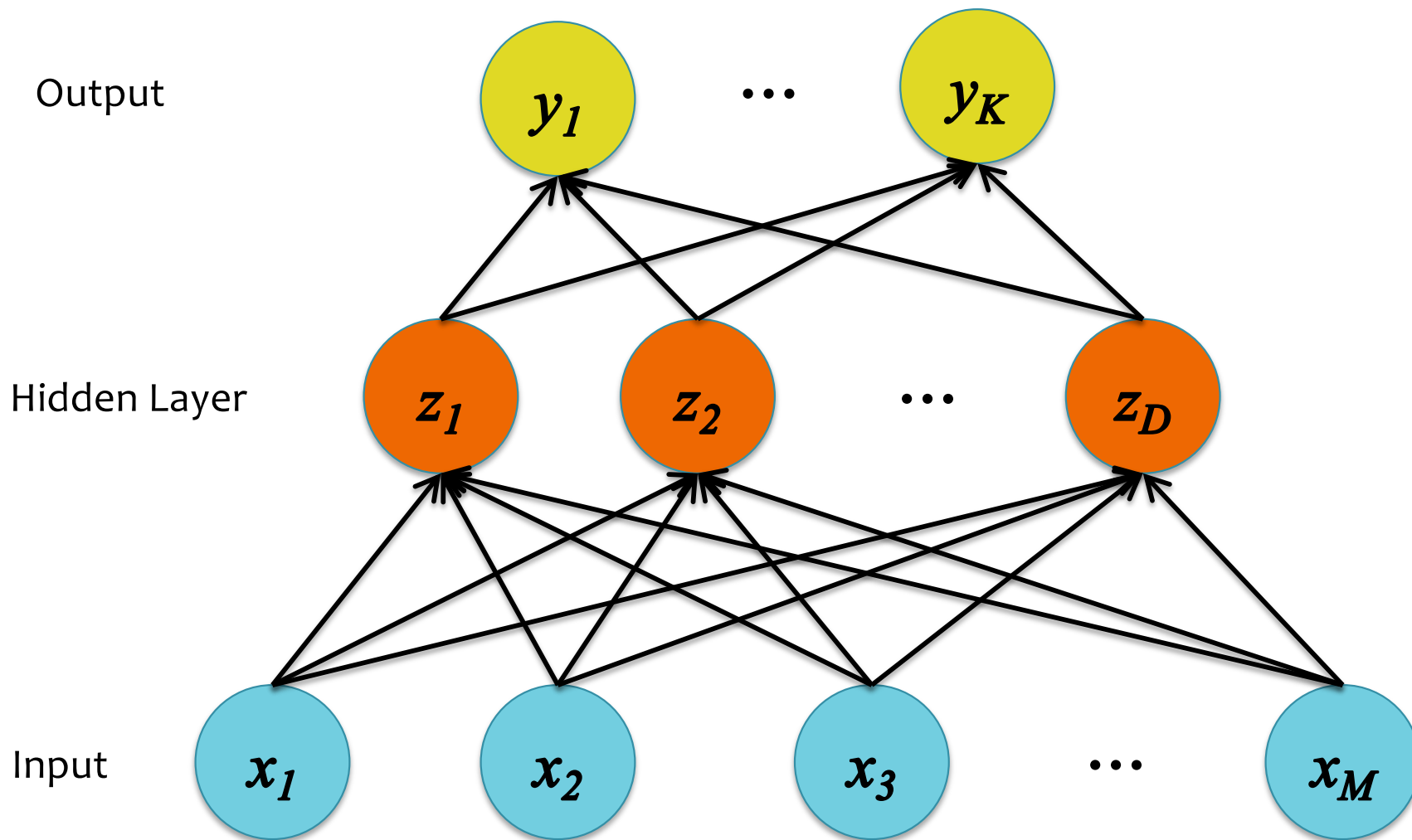


Figure 5: Cross entropy (black, surface on top) and quadratic (red, bottom surface) cost as a function of two weights (one at each layer) of a network with two layers, W_1 respectively on the first layer and W_2 on the second, output layer.

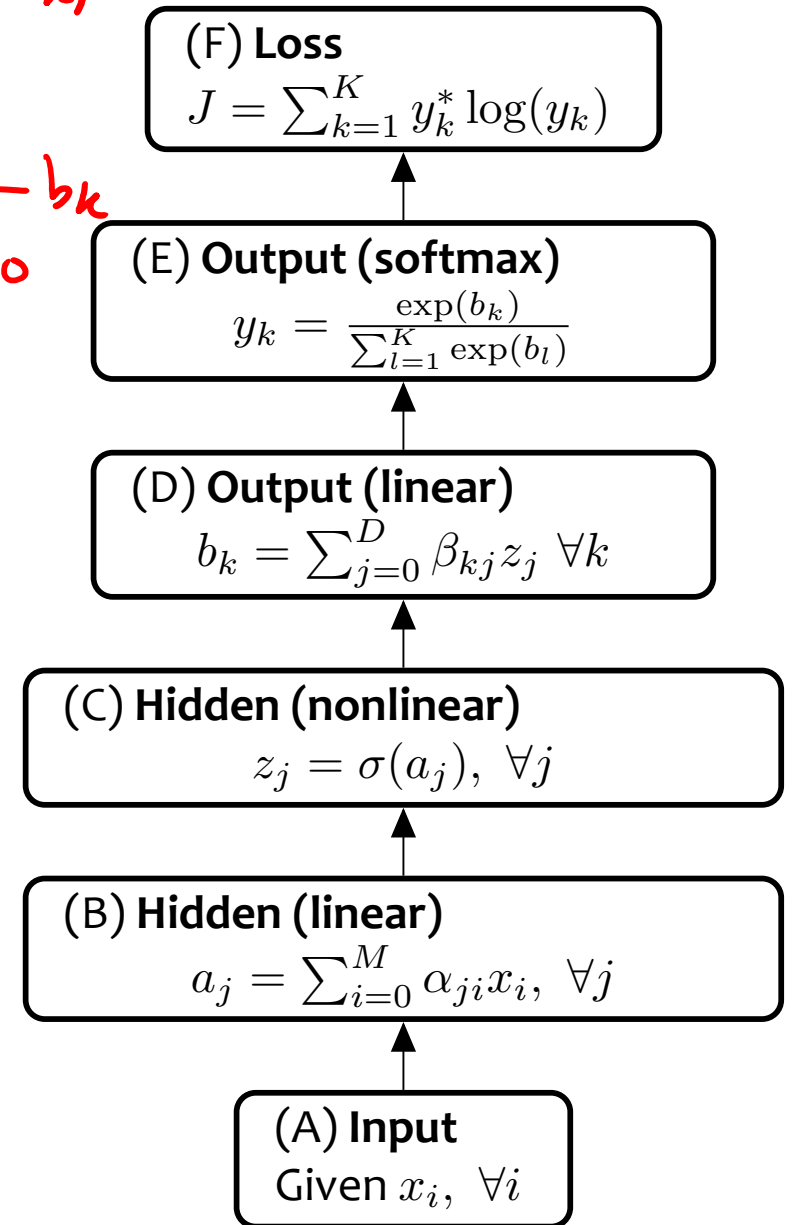
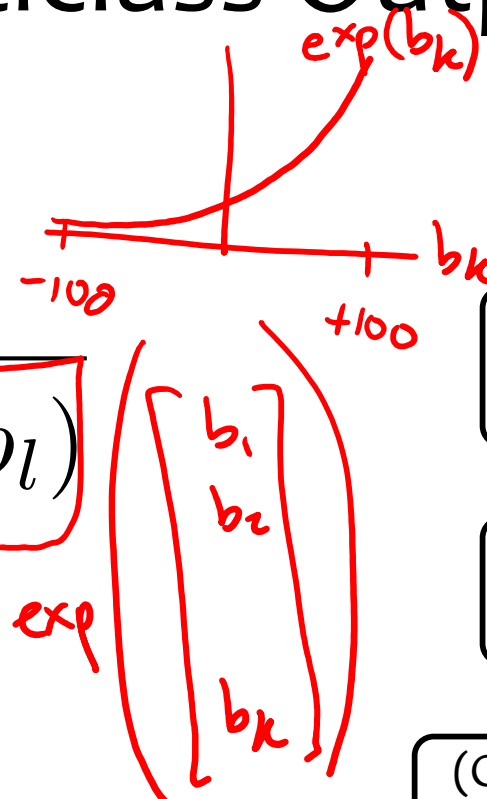
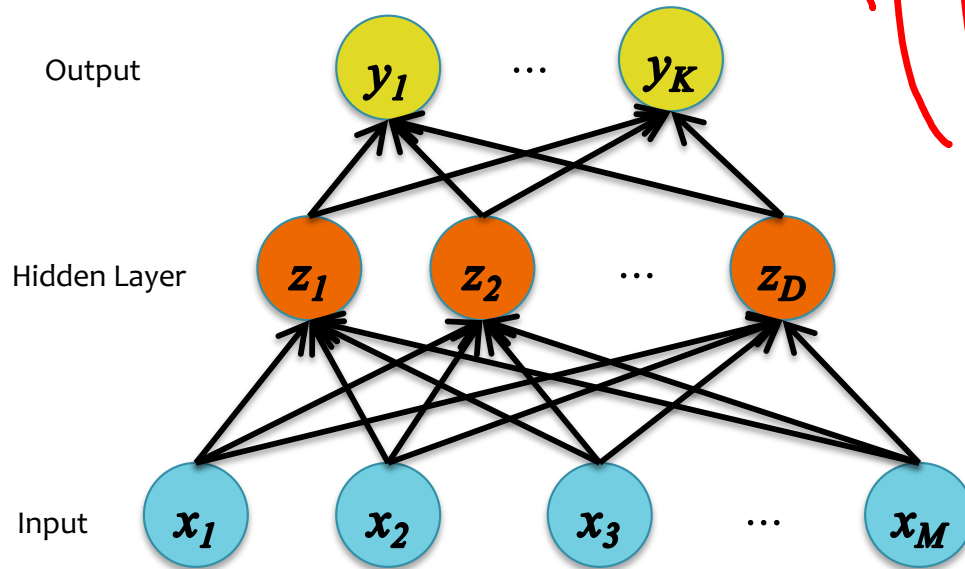
Multiclass Output



Multiclass Output

Softmax:

$$\underline{y_k} = \frac{\exp(b_k)}{\sum_{l=1}^K \exp(b_l)}$$



Objective Functions for NNs

3. Cross-Entropy for Multiclass Outputs:

- i.e. negative log likelihood for multiclass outputs
- Suppose output is a random variable Y that takes one of K values
- Let $\mathbf{y}^{(i)}$ represent our true label as a one-hot vector:

$$\mathbf{y}^{(i)} = \begin{array}{|c|c|c|c|c|c|c|c|} \hline 0 & 0 & 0 & 1 & 0 & 0 & \dots & 0 \\ \hline 1 & 2 & 3 & 4 & 5 & 6 & \dots & K \\ \hline \end{array}$$

true label = 4 = ℓ

- Assume our model outputs a length K vector of probabilities:

$$\mathbf{y} = \text{softmax}(f_{\text{scores}}(\mathbf{x}, \boldsymbol{\theta}))$$

- Then we can write the log-likelihood of a single training example $(\mathbf{x}^{(i)}, \mathbf{y}^{(i)})$ as:

$$J = \ell_{CE}(\mathbf{y}, \mathbf{y}^{(i)}) = - \sum_{k=1}^K y_k^{(i)} \log(y_k)$$

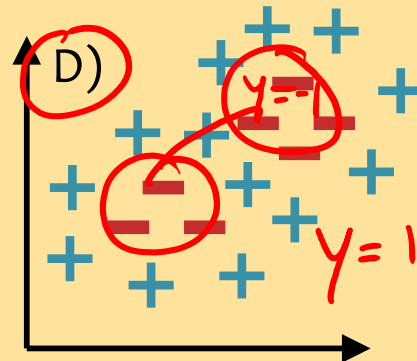
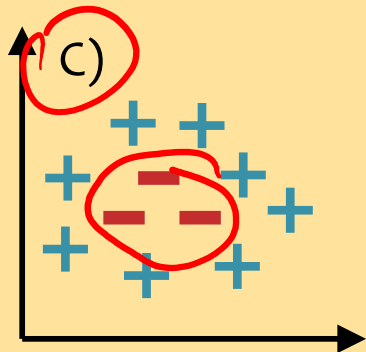
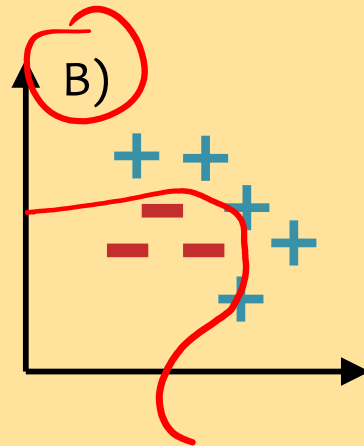
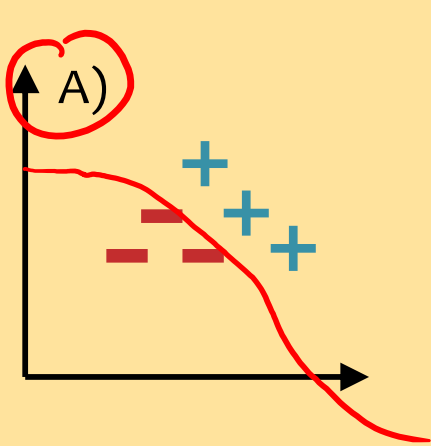
$$\begin{aligned} &= -\log(y_{\ell}) \\ &= -\log(\text{Pr}[Y = \ell]) \end{aligned}$$

Q₁

Neural Network Errors

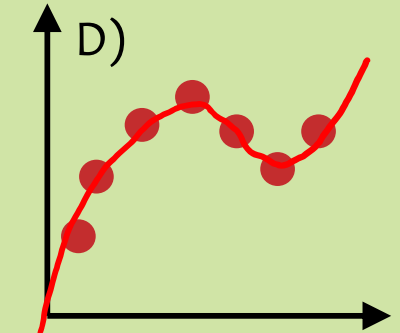
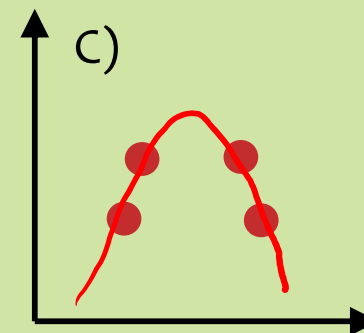
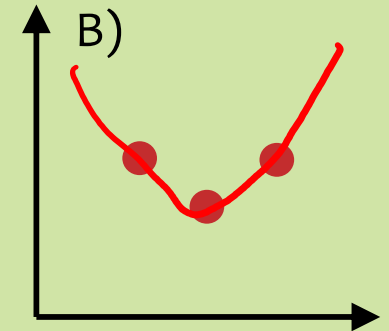
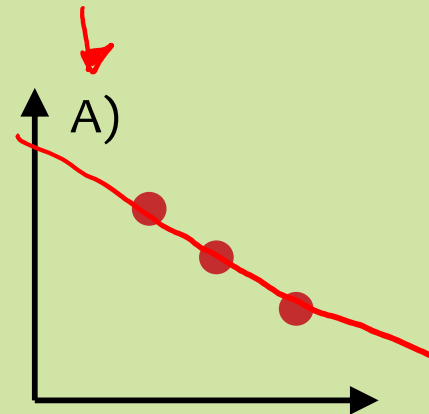
Q₂

Question X: For which of the datasets below does there exist a one-hidden layer neural network that achieves zero *classification* error? **Select all that apply.**



E = bxic

Question Y: For which of the datasets below does there exist a one-hidden layer neural network for *regression* that achieves *nearly zero MSE*? **Select all that apply.**



E = toxic

Q3: What questions do you have?

Neural Networks Objectives

You should be able to...

- Explain the biological motivations for a neural network
- Combine simpler models (e.g. linear regression, binary logistic regression, multinomial logistic regression) as components to build up feed-forward neural network architectures
- Explain the reasons why a neural network can model nonlinear decision boundaries for classification
- Compare and contrast feature engineering with learning features
- Identify (some of) the options available when designing the architecture of a neural network
- Implement a feed-forward neural network

Computing Gradients

APPROACHES TO DIFFERENTIATION

Background

A Recipe for Machine Learning

1. Given training data:

$$\{\mathbf{x}_i, \mathbf{y}_i\}_{i=1}^N$$

2. Choose each of these:

– Decision function

$$\hat{\mathbf{y}} = f_{\boldsymbol{\theta}}(\mathbf{x}_i)$$

– Loss function

$$\ell(\hat{\mathbf{y}}, \mathbf{y}_i) \in \mathbb{R}$$

3. Define goal:

$$\boldsymbol{\theta}^* = \arg \min_{\boldsymbol{\theta}} \sum_{i=1}^N \ell(f_{\boldsymbol{\theta}}(\mathbf{x}_i), \mathbf{y}_i)$$

4. Train with SGD:

(take small steps
opposite the gradient)

$$\boldsymbol{\theta}^{(t+1)} = \boldsymbol{\theta}^{(t)} - \eta_t \nabla \ell(f_{\boldsymbol{\theta}}(\mathbf{x}_i), \mathbf{y}_i)$$

Background

1. Given training data

$$\{\mathbf{x}_i, \mathbf{y}_i\}_{i=1}^N$$

2. Choose each of the

– Decision function

$$\hat{\mathbf{y}} = f_{\boldsymbol{\theta}}(\mathbf{x}_i)$$

– Loss function

$$\ell(\hat{\mathbf{y}}, \mathbf{y}_i) \in \mathbb{R}$$

Gradients

Backpropagation can compute this gradient!

And it's a **special case of a more general algorithm** called reverse-mode automatic differentiation that can compute the gradient of any differentiable function efficiently!

opposite the gradient)


$$\boldsymbol{\theta}^{(t+1)} = \boldsymbol{\theta}^{(t)} - \eta_t \nabla \ell(f_{\boldsymbol{\theta}}(\mathbf{x}_i), \mathbf{y}_i)$$

- **Question 1:**
When can we compute the gradients for an arbitrary neural network?
- **Question 2:**
When can we make the gradient computation efficient?

1. Finite Difference Method

- Pro: Great for testing implementations of backpropagation
- Con: Slow for high dimensional inputs / outputs
- Required: Ability to call the function $f(\mathbf{x})$ on any input $\mathbf{x}$

2. Symbolic Differentiation

- Note: The method you learned in high-school
- Note: Used by Mathematica / Wolfram Alpha / Maple
- Pro: Yields easily interpretable derivatives
- Con: Leads to exponential computation time if not carefully implemented
- Required: Mathematical expression that defines $f(\mathbf{x})$

Given $f : \mathbb{R}^A \rightarrow \mathbb{R}^B, f(\mathbf{x})$
Compute $\frac{\partial f(\mathbf{x})_i}{\partial x_j} \forall i, j$

Approaches to Differentiation

$$\begin{aligned} J(\vec{\theta}) &\in \mathbb{R} \\ \vec{\theta} &\in \mathbb{R}^M \end{aligned}$$

3. Automatic Differentiation - Reverse Mode
 - Note: Called *Backpropagation* when applied to Neural Nets
 - Pro: Computes partial derivatives of one output $f(\mathbf{x})_i$ with respect to all inputs x_j in time proportional to computation of $f(\mathbf{x})$
 - Con: Slow for high dimensional outputs (e.g. vector-valued functions)
 - Required: Algorithm for computing $f(\mathbf{x})$
4. Automatic Differentiation - Forward Mode
 - Note: Easy to implement. Uses dual numbers.
 - Pro: Computes partial derivatives of all outputs $f(\mathbf{x})_i$ with respect to one input x_j in time proportional to computation of $f(\mathbf{x})$
 - Con: Slow for high dimensional inputs (e.g. vector-valued $\mathbf{x}$)
 - Required: Algorithm for computing $f(\mathbf{x})$

Given $f : \mathbb{R}^A \rightarrow \mathbb{R}^B$, $f(\mathbf{x})$
Compute $\frac{\partial f(\mathbf{x})_i}{\partial x_j} \forall i, j$

THE FINITE DIFFERENCE METHOD

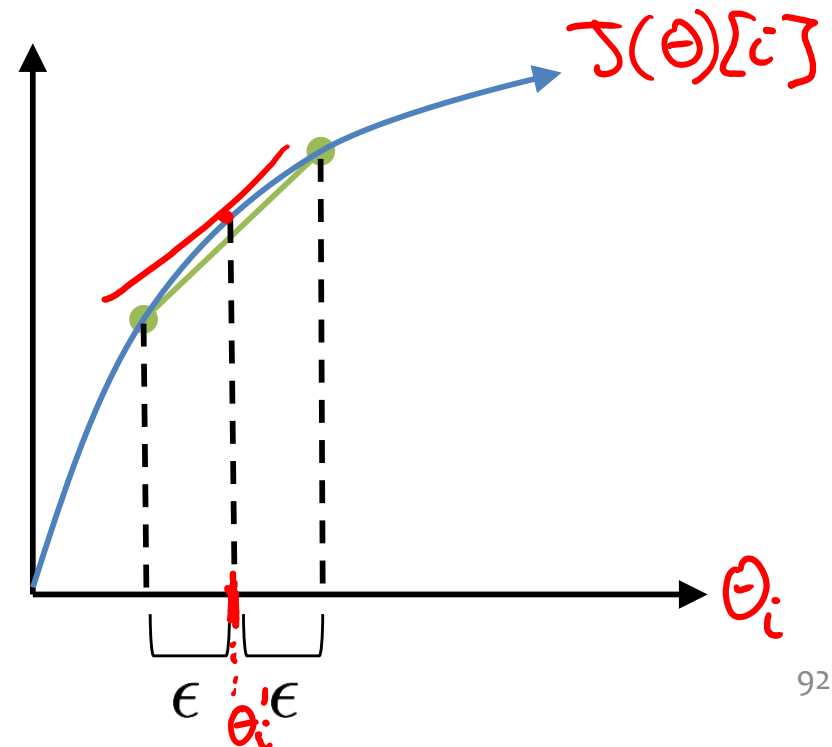
The *centered* finite difference approximation is:

$$\frac{\partial}{\partial \theta_i} J(\boldsymbol{\theta}) \approx \frac{(J(\boldsymbol{\theta} + \epsilon \cdot \mathbf{d}_i) - J(\boldsymbol{\theta} - \epsilon \cdot \mathbf{d}_i))}{2\epsilon} \quad (1)$$

where $\mathbf{d}_i$ is a 1-hot vector consisting of all zeros except for the i th entry of $\mathbf{d}_i$, which has value 1.

Notes:

- Suffers from issues of floating point precision, in practice
- Typically only appropriate to use on small examples with an appropriately chosen epsilon



Speed Quiz:
2 minute time limit.

Differentiation Quiz #1:

Suppose $x = 2$ and $z = 3$, what are dy/dx and dy/dz for the function below? **Round your answer to the nearest integer.**

$$y = \exp(xz) + \frac{xz}{\log(x)} + \frac{\sin(\log(x))}{xz}$$

Answer: Answers below are in the form $[dy/dx, dy/dz]$

- | | |
|---------------|----------------|
| A. [42, -72] | E. [1208, 810] |
| B. [72, -42] | F. [810, 1208] |
| C. [100, 127] | G. [1505, 94] |
| D. [127, 100] | H. [94, 1505] |

Differentiation Quiz #2:

A neural network with 2 hidden layers can be written as:

$$y = \sigma(\beta^T \sigma((\alpha^{(2)})^T \sigma((\alpha^{(1)})^T \mathbf{x}))$$

where $y \in \mathbb{R}$, $\mathbf{x} \in \mathbb{R}^{D^{(0)}}$, $\beta \in \mathbb{R}^{D^{(2)}}$ and $\alpha^{(i)}$ is a $D^{(i)} \times D^{(i-1)}$ matrix. Nonlinear functions are applied elementwise:

$$\sigma(\mathbf{a}) = [\sigma(a_1), \dots, \sigma(a_K)]^T$$

Let σ be sigmoid: $\sigma(a) = \frac{1}{1+\exp(-a)}$

What is $\frac{\partial y}{\partial \beta_j}$ and $\frac{\partial y}{\partial \alpha_j^{(i)}}$ for all i, j .

