

Foundations of Neural Networks

Daniel B. Rowe, PhD

Professor of Computational Statistics

Department of Mathematical and Statistical Sciences

Marquette University



Department of
Biophysics



March 24, 2023

School of Mathematical
& Statistical Sciences



Outline

1. Introduction

NN Structure, Activation/Score Functions, Estimation

2. Linear Regression and Neural Nets

Simple & Multivariate with Gradient Descent

3. Non-Linear Regression and Neural Nets

Simple & Multivariate with Gradient Descent

4. Logistic Regression and Neural Nets

Simple & Multivariate with Gradient Descent

5. Multi-Layer Deep Neural Nets

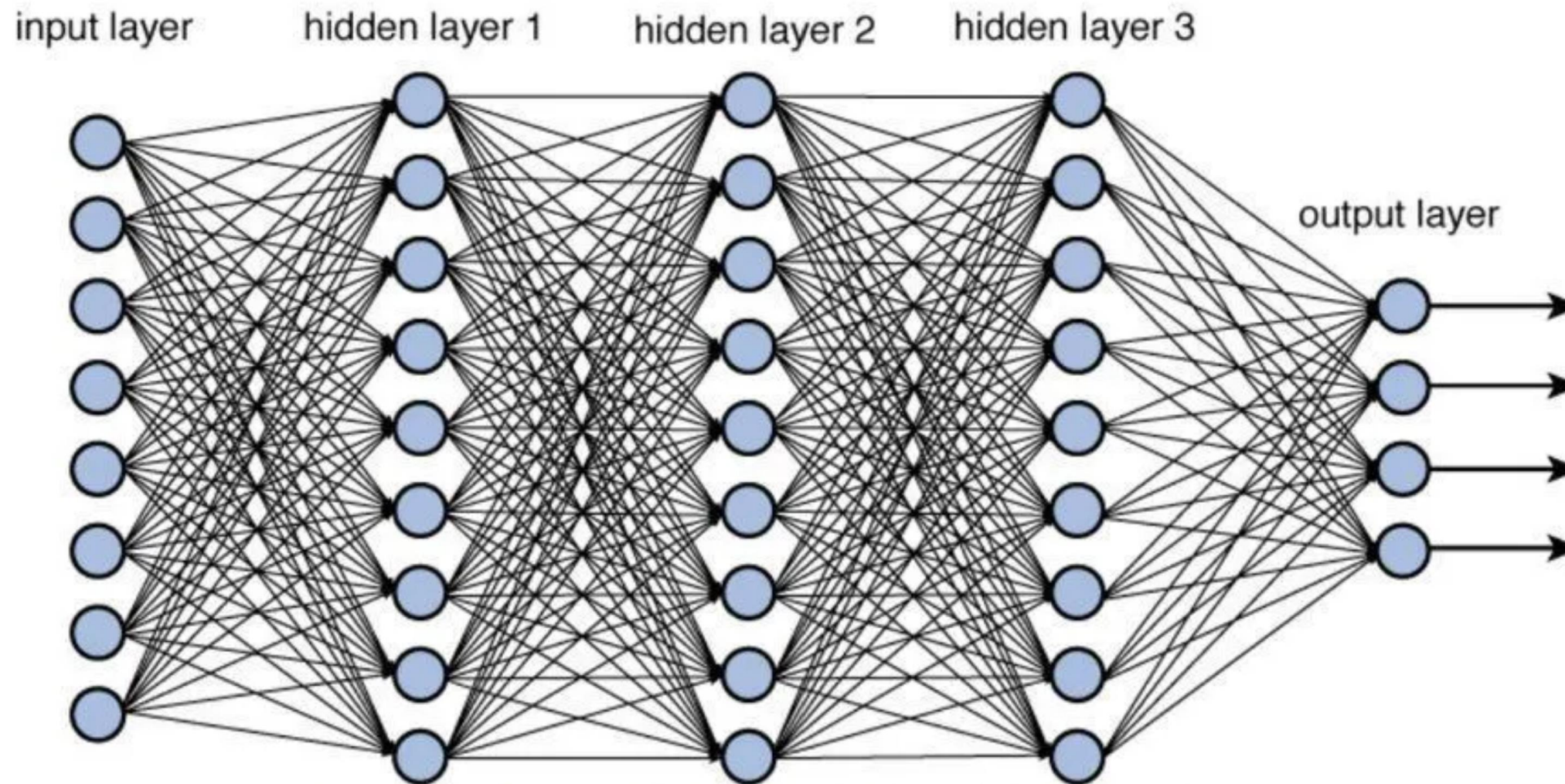
Two or More Layers

6. Discussion

More to learn hands on.

1. Introduction

There are often illustrations, but no details of mathematics.

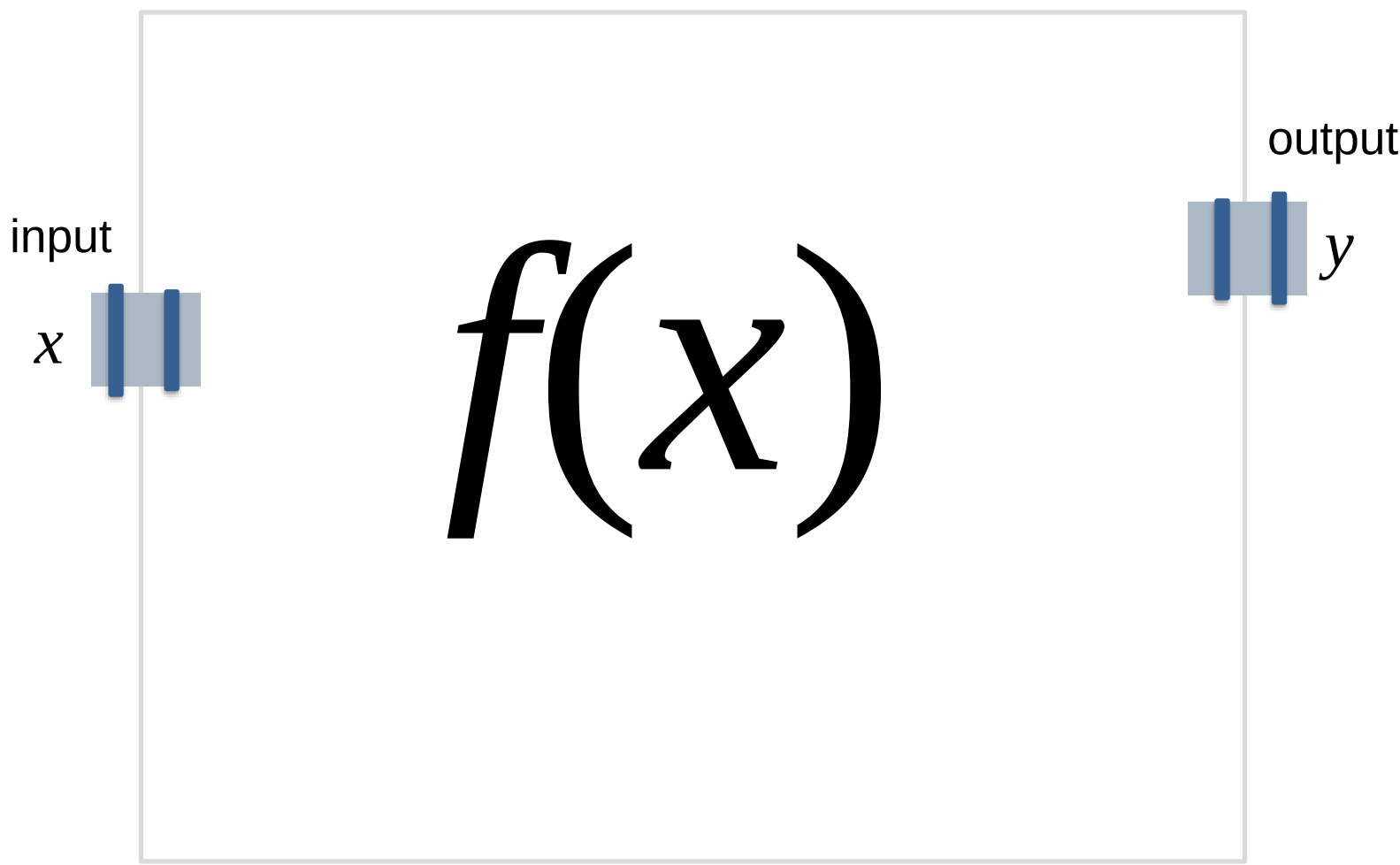


This leaves the science out of Data Science and results in a Data Artist.

<https://towardsdatascience.com/training-deep-neural-networks-9fdb1964b964>

1. Introduction

Assume we want to learn the relationship between x and y .



Sample	
outputs	inputs
y_1	x_1
y_2	x_2
y_3	x_3
y_4	x_4
y_5	x_5
y_6	x_6
y_7	x_7
y_8	x_8
y_9	x_9
y_{10}	x_{10}

1. Introduction

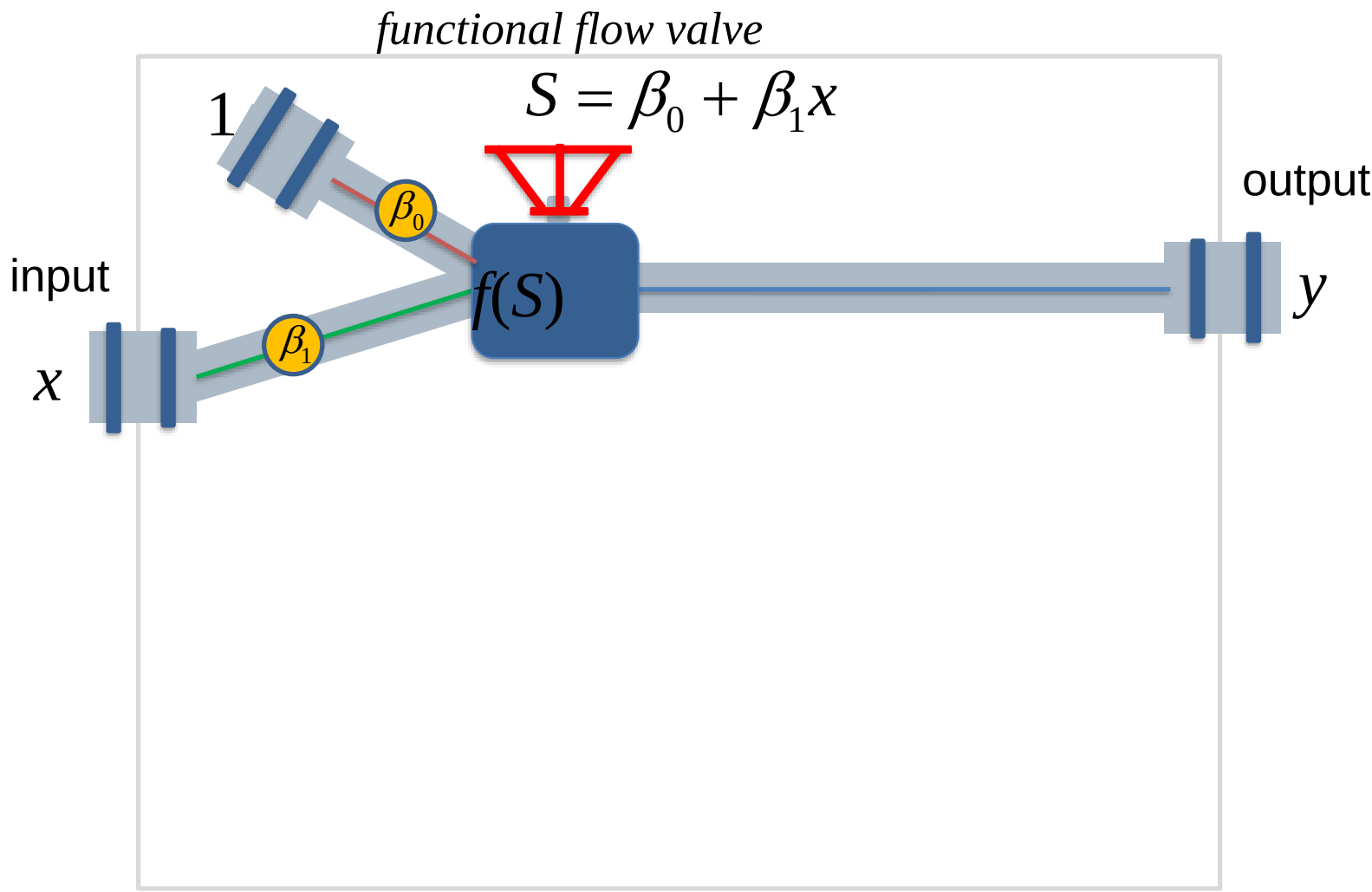
coefficient
 $j = 0, \dots, q$

β_j

$f(S)$

$p(S) = \frac{1}{1 + e^{-S}}$

Single input and single output.



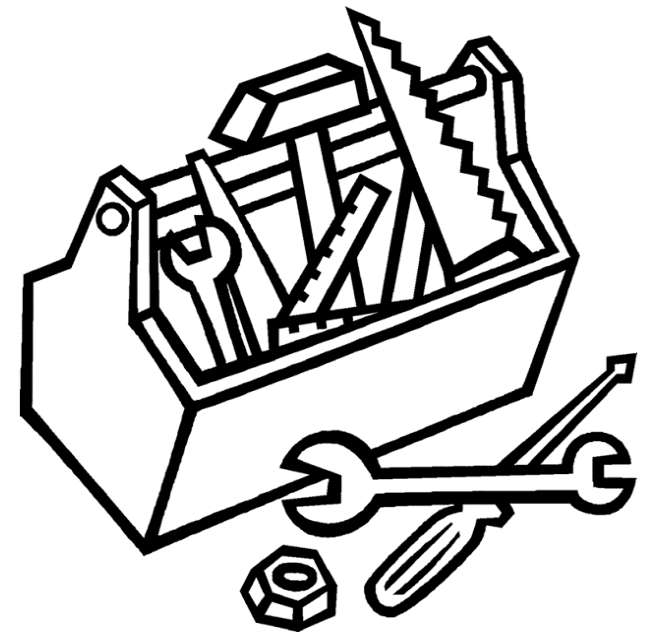
Functional Form

$y = S$

or

$y = \begin{cases} 1 & \text{w/ probability } p(S) \\ 0 & \text{w/ probability } 1-p(S) \end{cases}$

Many Tools



Well known solutions and interpretations.

1. Introduction

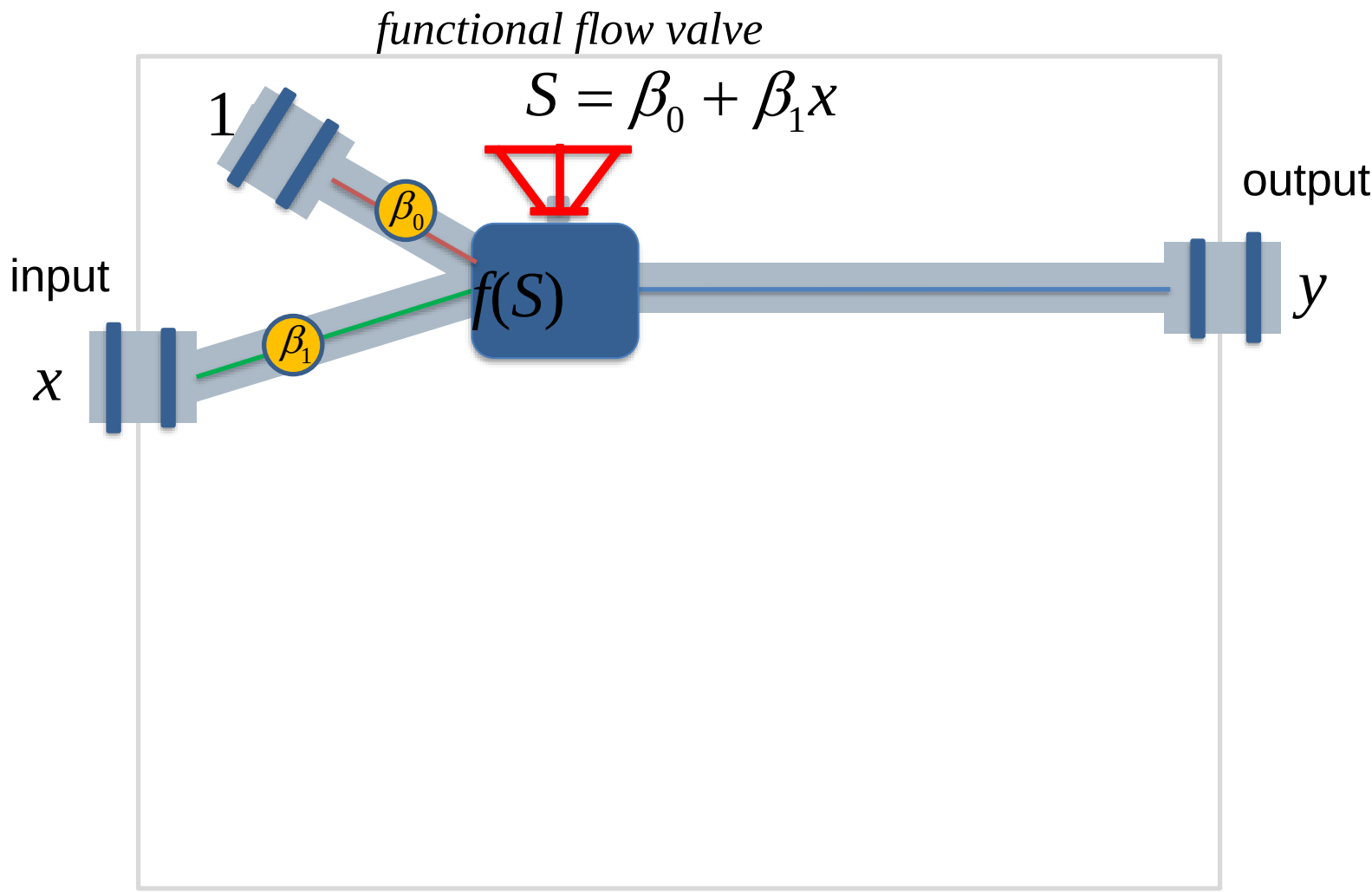
coefficient
 $j = 0, \dots, q$

β_j

$f(S)$

$p(S) = \frac{1}{1 + e^{-S}}$

Single input and single output.



Functional Form

$y = S$

or

$y = \begin{cases} 1 & \text{w/ probability } p(S) \\ 0 & \text{w/ probability } 1-p(S) \end{cases}$

Objective Function

$$Q(\beta_0, \dots, \beta_q)$$

Estimate Parameters

$$(\hat{\beta}_0, \hat{\beta}_1)$$

Well known solutions and interpretations.

1. Introduction

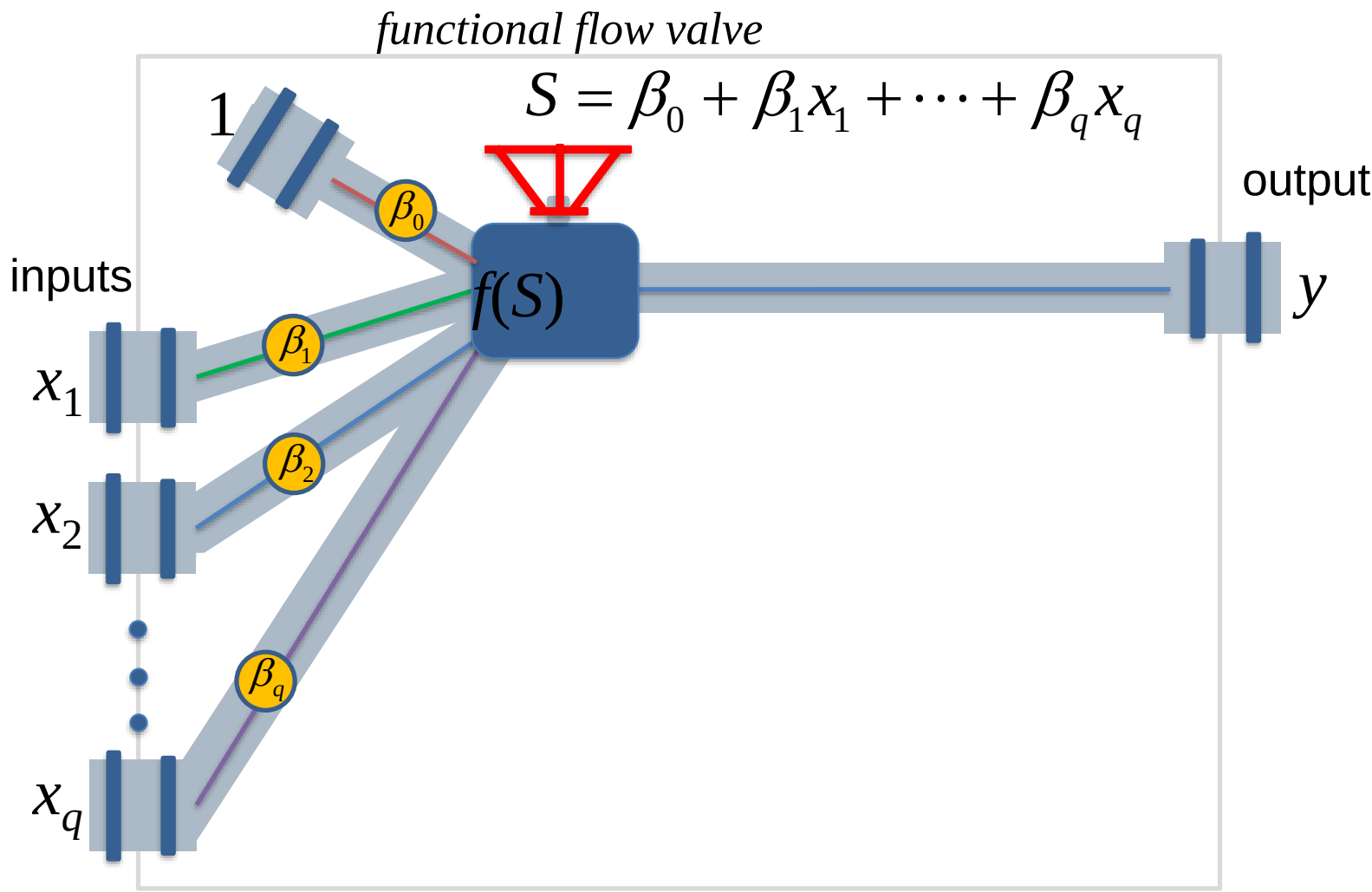
coefficient
 $j = 0, \dots, q$

β_j

$f(S)$

$p(S) = \frac{1}{1 + e^{-S}}$

Multiple inputs and single output.



Functional Form

$y = S$

or

$y = \begin{cases} 1 & \text{w/ probability } p(S) \\ 0 & \text{w/ probability } 1-p(S) \end{cases}$

Objective Function

$$Q(\beta_0, \dots, \beta_q)$$

Estimate Parameters

$$(\hat{\beta}_0, \dots, \hat{\beta}_q)$$

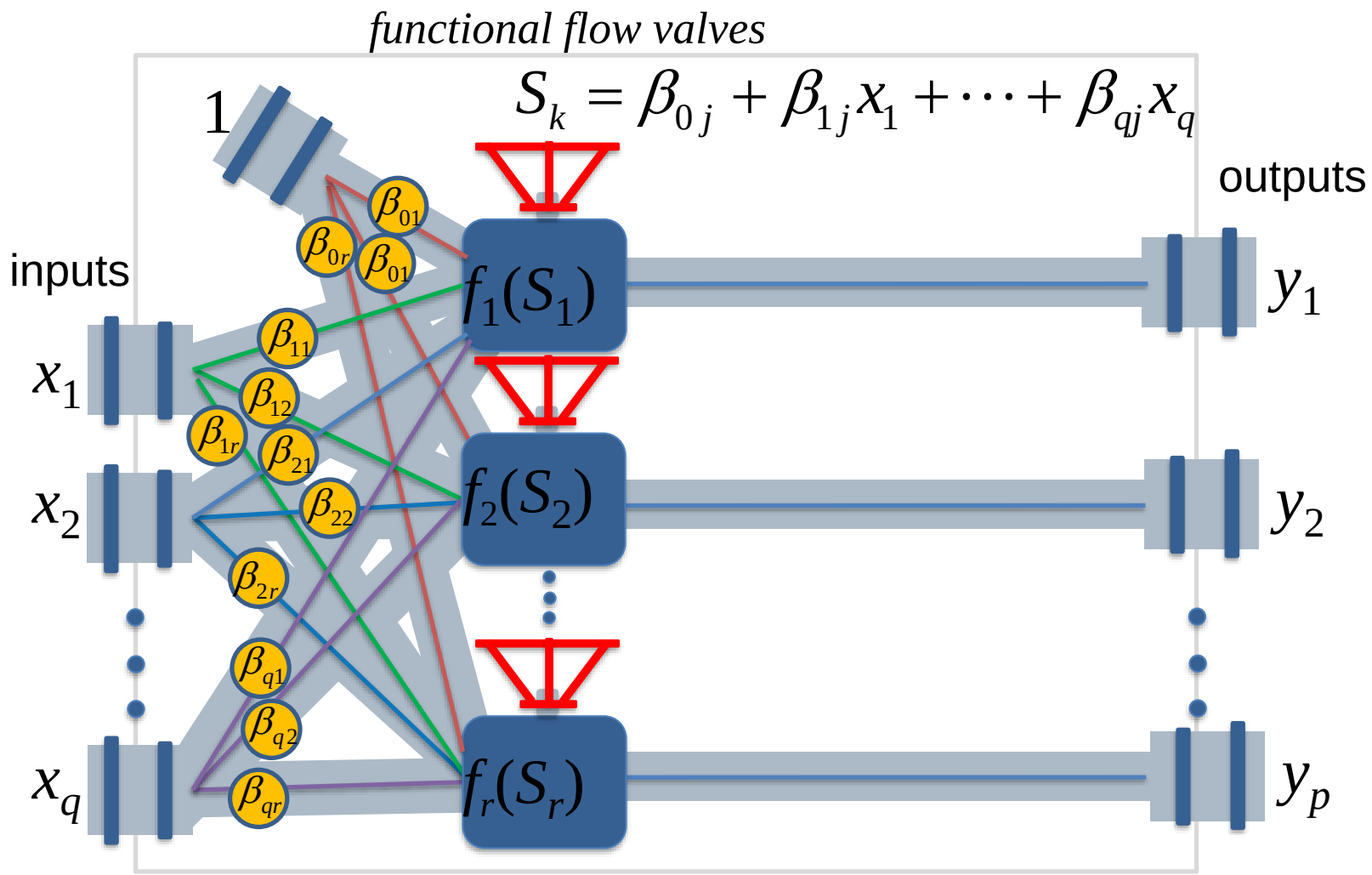
Well known solutions and interpretations.

1. Introduction

coefficient node
 $j = 0, \dots, q$ $k = 1, \dots, r$

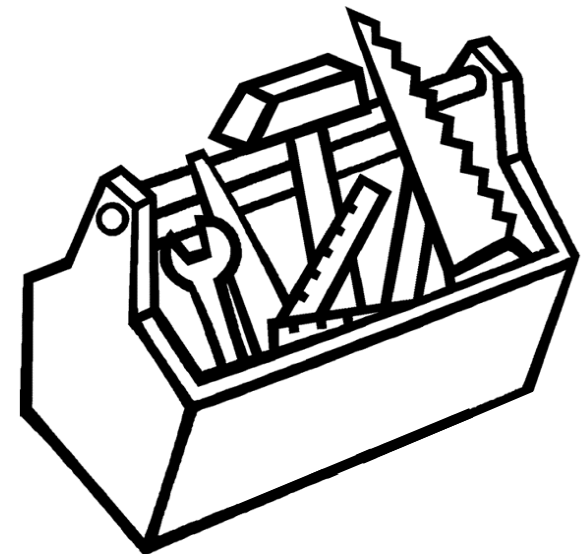
β_{jk} $f_k(S_k)$

Multiple inputs and multiple outputs.



Functional Form
Multivariate Linear Regression
Other nonlinear functions
Multinomial Logistic...

Fewer Tools



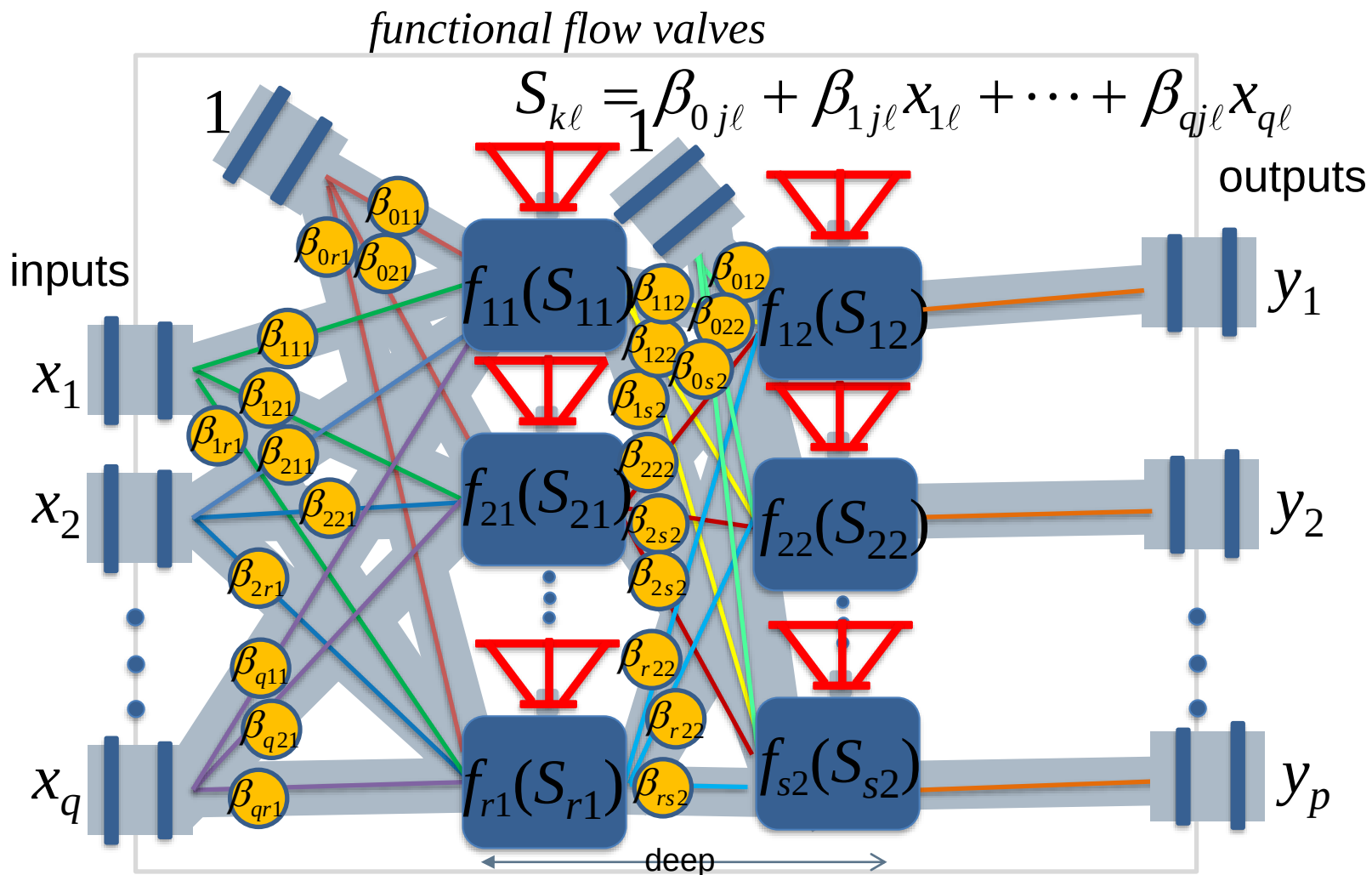
Well known solutions and interpretations.

1. Introduction

coefficient node layer
 $j = 0, \dots, q$ $k = 1, \dots, r$ $\ell = 1, 2$

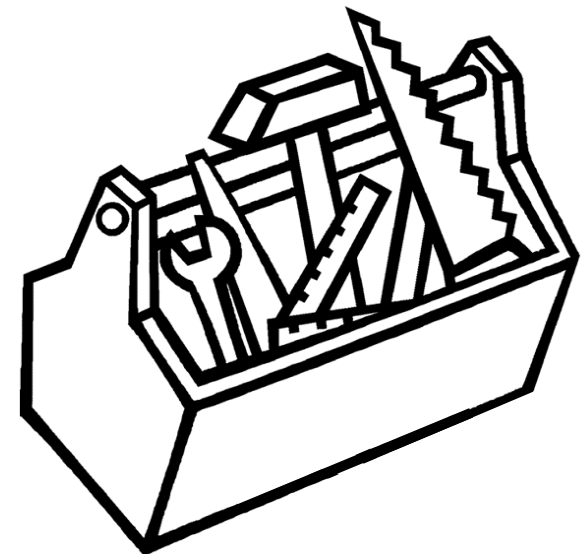
$$\beta_{jkl} \quad f_{kl}(S_{kl})$$

Multiple inputs and multiple outputs.



Functional Form
Multivariate Linear Regression
Other nonlinear functions
Multinomial Logistic...

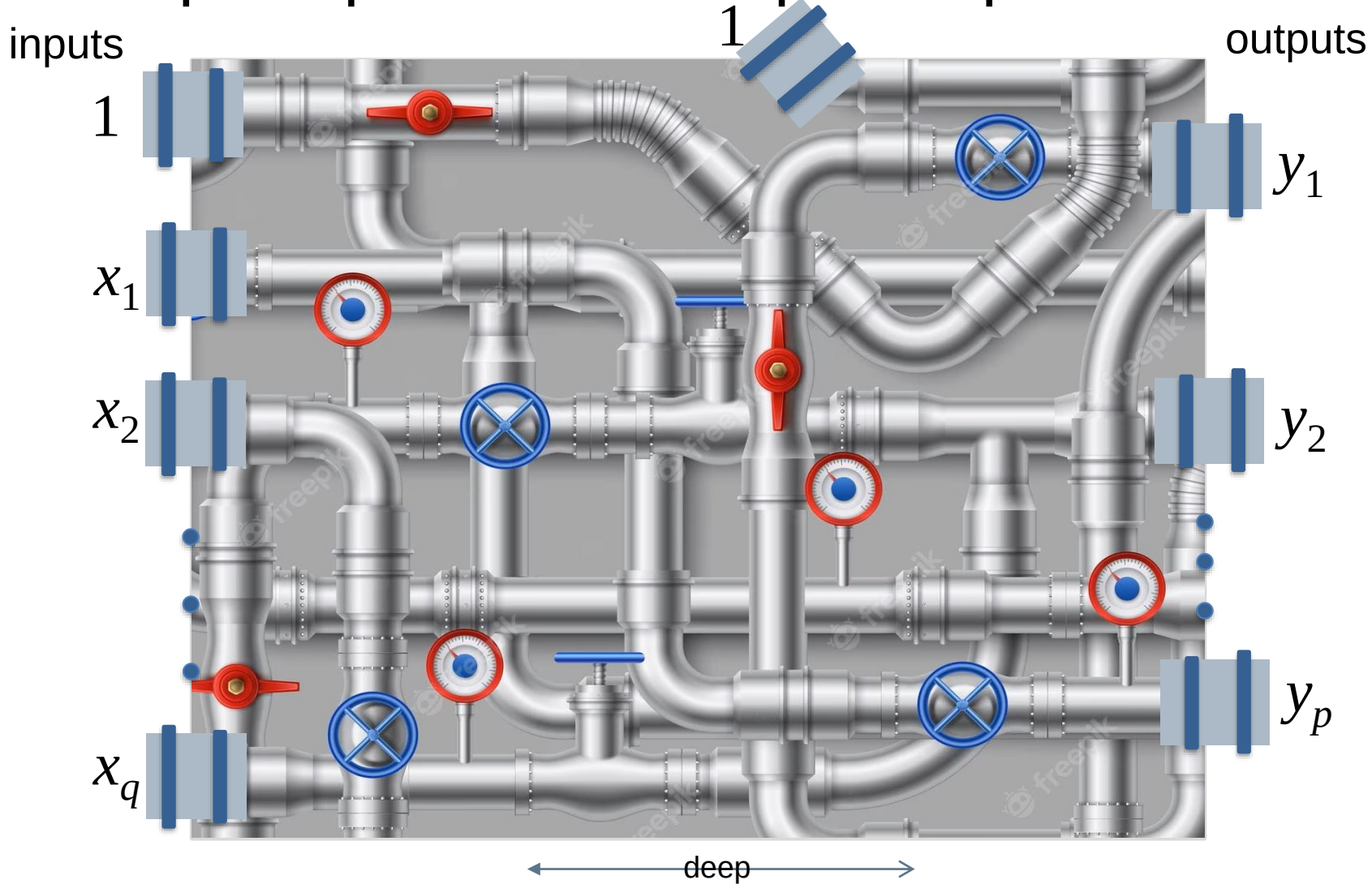
Fewer Tools



Well known solutions and interpretations.
Many interconnected valves. Functions of functions of ... of inputs.

1. Introduction

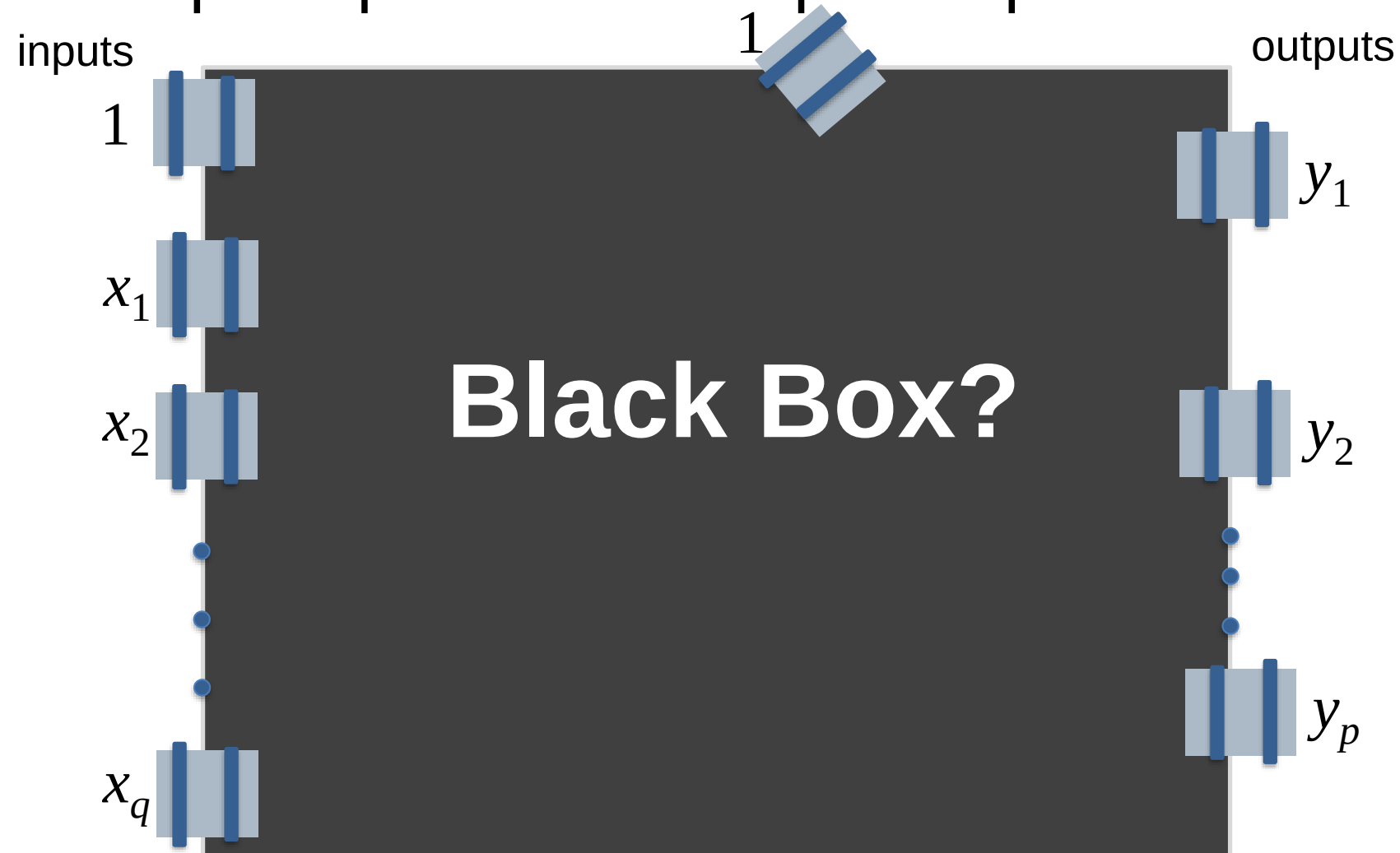
Multiple inputs and multiple output.



Many interconnected valves. Functions of functions of ... of inputs.

1. Introduction

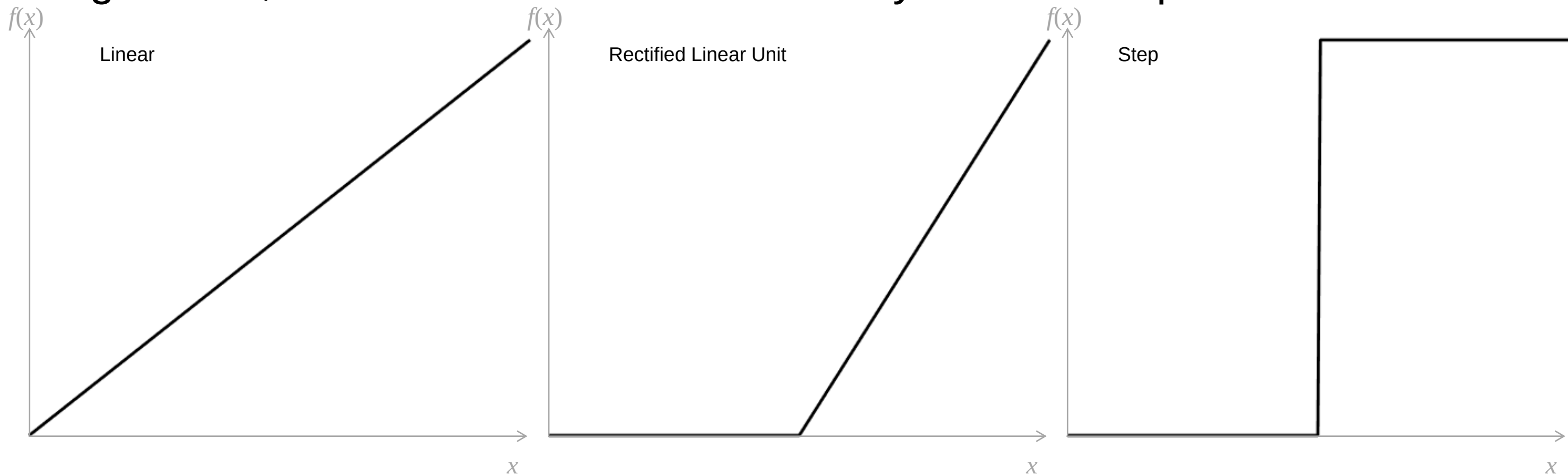
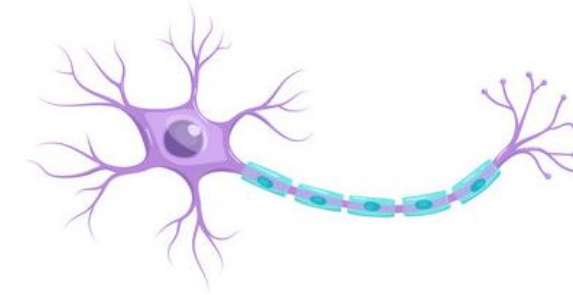
Multiple inputs and multiple output.



Although you put all the pieces in the box, it is extremely complicated to understand and interpret. We will examine the basic structure.

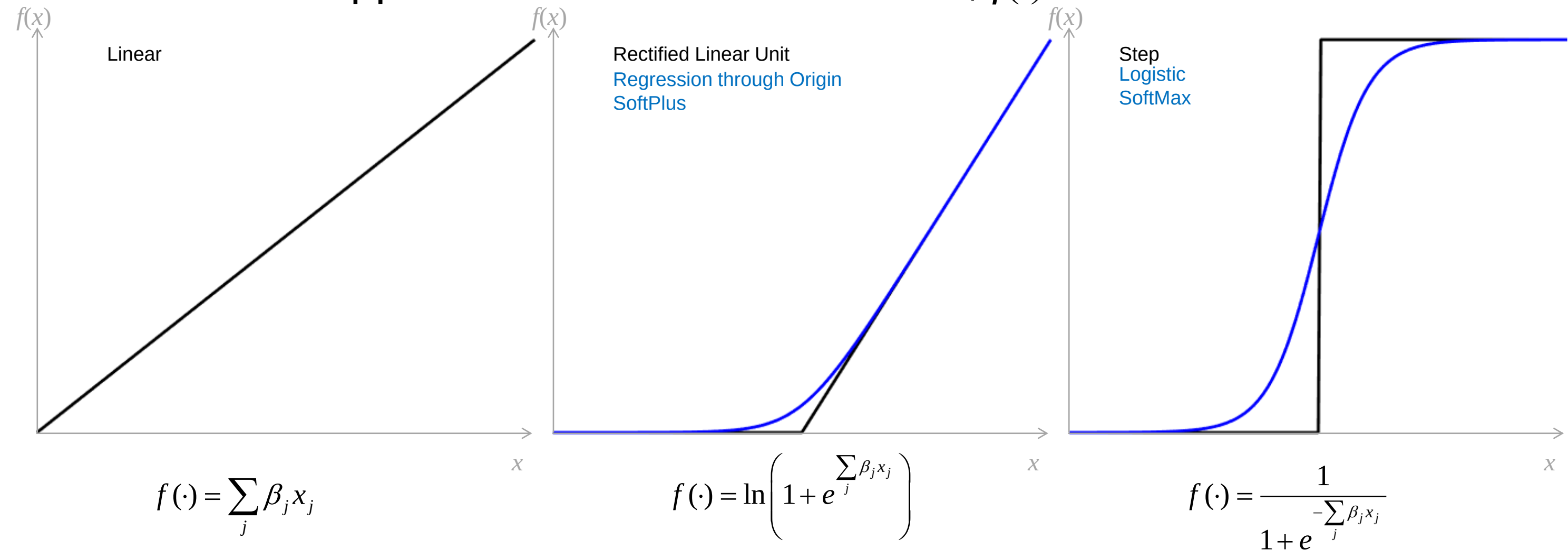
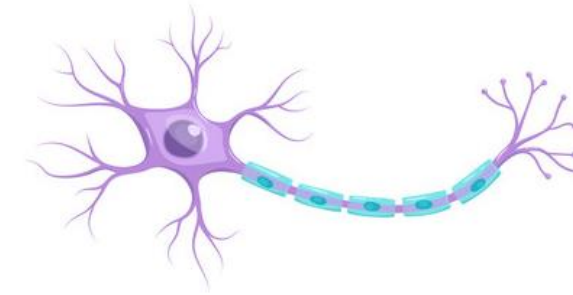
1. Introduction

There are many activation functions $f(\cdot)$ from (non)linear regression, but most used are motivated by neuronal representations.



1. Introduction

Step and ReLU not differentiable for optimization so smooth differentiable approximate activation functions, $f(\cdot)$ are used.



1. Introduction

$$S_i = \beta_0 + \beta_1 x_{1i} + \dots + \beta_q x_{qi}$$

The form of the function depends on the type of input/output variable.

Distribution	Activation	Likelihood
Continuous Real Valued $y_i \sim Normal(f(S_i), \sigma^2)$	Linear $f(S_i) = S_i$	Normal $L(\beta_1, \dots, \beta_q) = (2\pi)^{-n/2} \exp\left[-\frac{1}{2\sigma^2} \sum_i (y_i - f(S_i))^2\right]$
Continuous Real Valued (Not focusing on.) $y_i \sim Normal(f(S_i), \sigma^2)$	SoftPlus $f(S_i) = \ln(1 + e^{S_i})$	Normal $L(\beta_1, \dots, \beta_q) = (2\pi)^{-n/2} \exp\left[-\frac{1}{2\sigma^2} \sum_i (y_i - f(S_i))^2\right]$
Discrete 0/1 Valued $y_i \sim Bernoulli(p(S_i))$	Logistic $p(S_i) = \frac{1}{1 + e^{-S_i}}$	Bernoulli $L(\beta_0, \dots, \beta_q) = \prod_i [p(S_i)]^{y_i} [1 - p(S_i)]^{1-y_i}$

These are motivated from statistics!

1. Introduction

$$S_i = \beta_0 + \beta_1 x_{1i} + \dots + \beta_q x_{qi}$$
$$= \sum_j \beta_j x_{ji}$$

We generally take the natural log of the likelihood for the score function.

Distribution	Activation	Score/Objective Function
Continuous Real Valued	Linear	Least Squares
$y_i \sim Normal(f(S_i), \sigma^2)$	$f(S_i) = S_i$	$Q = \frac{1}{n} \sum_i [y_i - \sum_j \beta_j x_{ji}]^2$
Continuous Real Valued (Not focusing on.)	SoftPlus	Least Squares
$y_i \sim Normal(f(S_i), \sigma^2)$	$f(S_i) = \ln(1 + e^{S_i})$	$Q = \frac{1}{n} \sum_i [y_i - \ln(1 + \exp(\sum_j \beta_j x_{ji}))]^2$
Discrete 0/1 Valued	Logistic	Bernoulli
$y_i \sim Bernoulli(p(S_i))$	$p(S_i) = \frac{1}{1 + e^{-S_i}}$	$Q = \sum_i y_i (\sum_j \beta_j x_{ji}) - \sum_i \ln[1 + \exp(\sum_j \beta_j x_{ji})]$

And optimize the score function using “training” data $x_1, \dots, x_n$!

1. Introduction

$$S_i = \beta_0 + \beta_1 x_{1i} + \cdots + \beta_q x_{qi} \\ = \sum_j \beta_j x_{ji}$$

We can differentiate the score function for optimization.

Distribution

Activation

Score/Objective Function

Least Squares-Linear

$$Q = \frac{1}{n} \sum_i [y_i - \sum_j \beta_j x_{ji}]^2$$

Derivative

$$\frac{\partial Q}{\partial \beta_j} = \frac{2}{n} \sum_i [y_i - \sum_j \beta_j x_{ij}] (-x_{ij})$$

Least Squares-SoftPlus
(Not focusing on.)

$$Q = \frac{1}{n} \sum_i [y_i - \ln(1 + \exp(\sum_j \beta_j x_{ji}))]^2$$

Derivative

$$\frac{\partial Q}{\partial \beta_j} = \frac{2}{n} \sum_i \left[\frac{x_{ij} [y_i - \ln(1 + \exp(\sum_j \beta_j x_{ij}))] \exp(\sum_j \beta_j x_{ij})}{1 + \exp(\sum_j \beta_j x_{ij})} \right]$$

Bernoulli-Logistic

$$Q = \sum_i y_i (\sum_j \beta_j x_{ji}) - \sum_i \ln[1 + \exp(\sum_j \beta_j x_{ji})]$$

Derivative

$$\frac{\partial Q}{\partial \beta_j} = \sum_i x_{ij} y_i - \sum_i \frac{x_{ij}}{1 + \exp(-\sum_j \beta_j x_{ij})}$$

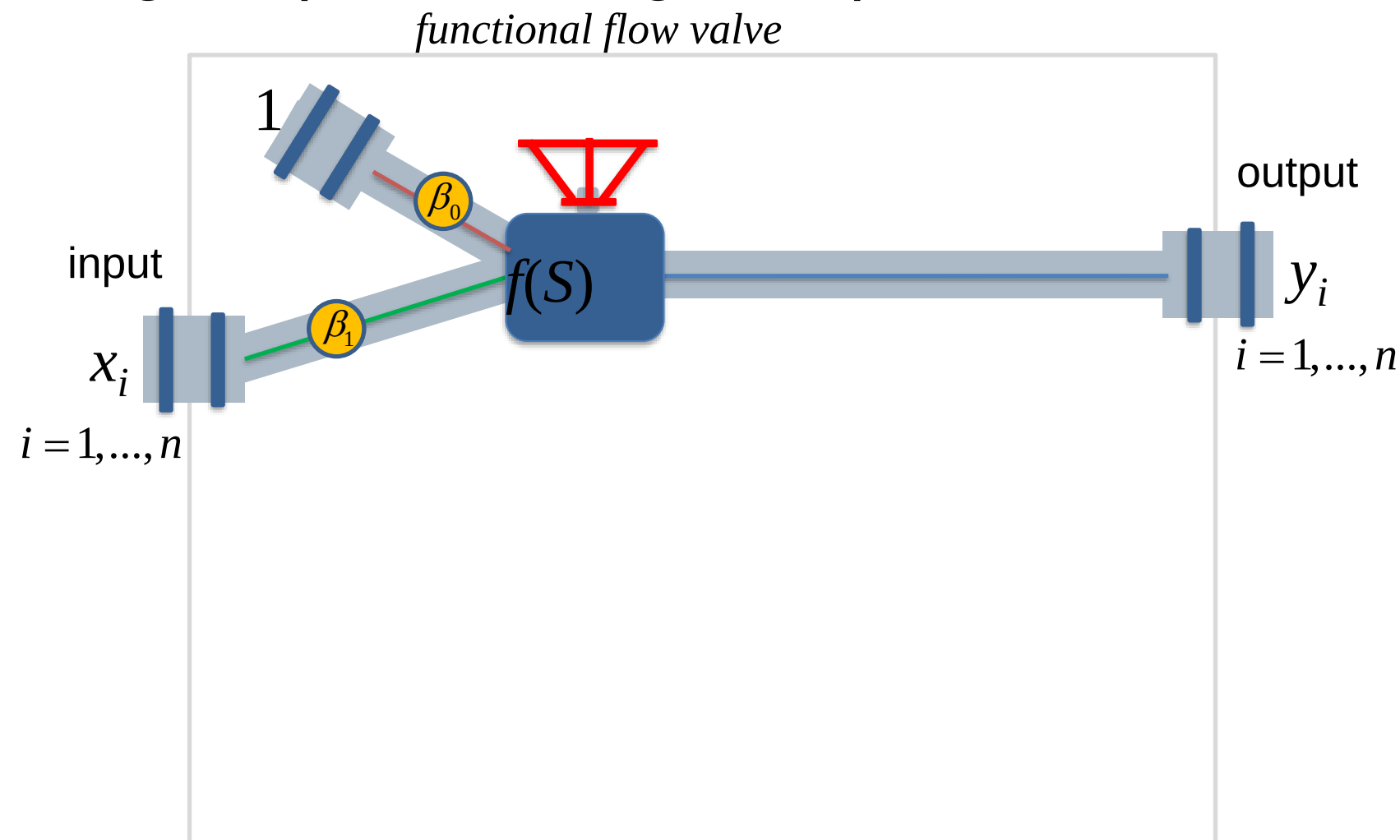
And optimize the score function using “training” data $x_1, \dots, x_n$!

1. Introduction

coefficient
 $j = 0, \dots, q$

$$\beta_j \quad f(S) \quad S = \sum_j \beta_j x_{ji}$$

Single input and single output.



A well known general numerical solution.

1) Start with initial $t=0$ values $(\hat{\beta}_0^{(t)}, \hat{\beta}_1^{(t)})$

2) Run n data through

$$Q_i^{(t)} = [y_i - f(\sum_j \hat{\beta}_j^{(t)} x_{ij})]^2$$

$i = 1, \dots, n$

3) Calculate score function

$$Q^{(t)} = \frac{1}{n} \sum_i [y_i - f(\sum_j \hat{\beta}_j^{(t)} x_{ij})]^2$$

4) Update coefficients GD

$$(\hat{\beta}_0^{(t+1)}, \hat{\beta}_1^{(t+1)}) \quad \nabla Q$$

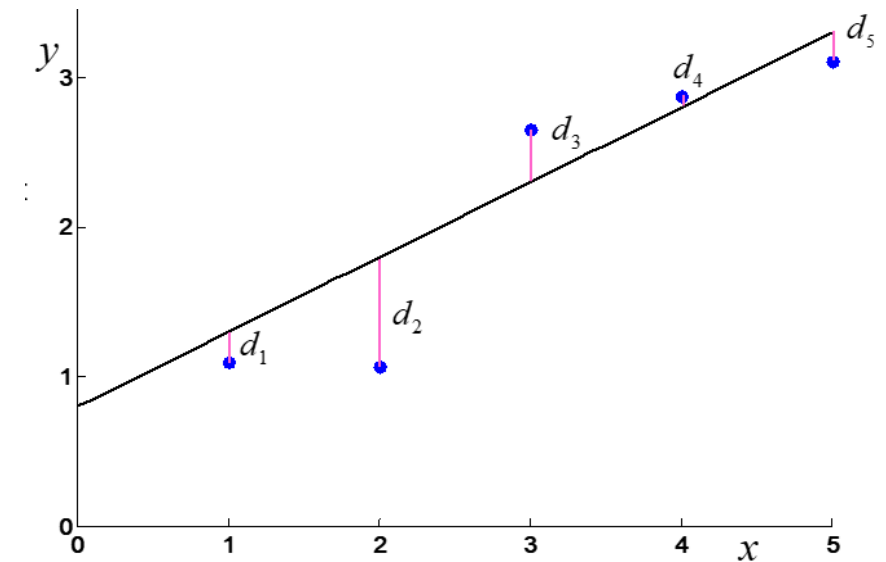
5) Return to Step 2.

2. Linear Regression and Neural Nets

Often we believe that there is a linear relationship between an independent variable x and a dependent variable y with measurement error.

$$y_i = \beta_0 + \beta_1 x_i + \varepsilon_i$$

$$i = 1, \dots, n$$



Could assume normal error or use least squares.

$$Q = \frac{1}{n} \sum_{i=1}^n (y_i - \beta_0 - \beta_1 x_i)^2$$

2. Linear Regression and Neural Nets

$$y_i = \beta_0 + \beta_1 x_i + \varepsilon_i$$

$$i = 1, \dots, n$$

We can estimate the “best” linear relationship between an independent variable x and a dependent variable y using the score function

$$Q = \frac{1}{n} \sum_{i=1}^n (y_i - \beta_0 - \beta_1 x_i)^2$$

by taking derivatives with respect to the β coefficients:

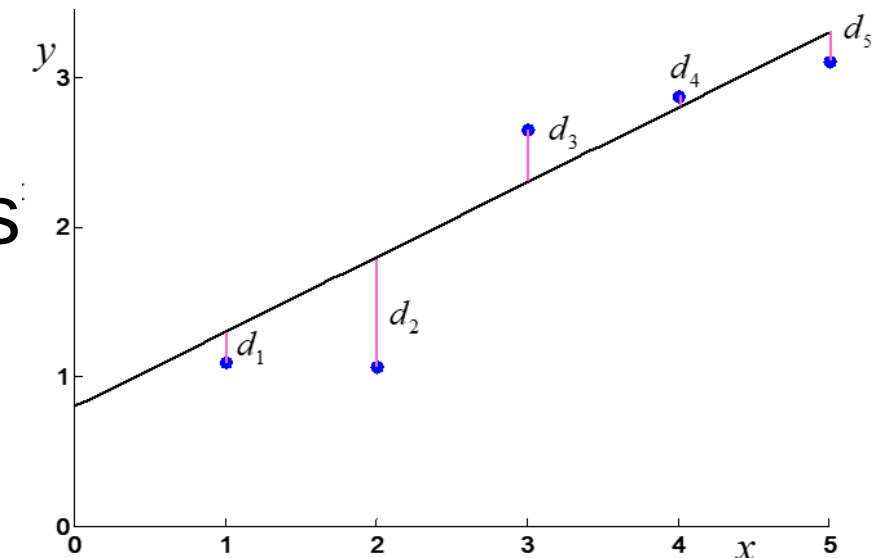
$$\frac{\partial Q}{\partial \beta_j} = -\frac{2}{n} \sum_i x_{ij} (y_i - \sum_j \beta_j x_{ij}) \quad x_{i0} = 1 \quad \begin{matrix} i = 1, \dots, n \\ j = 1, \dots, q \end{matrix}$$

written in vector form as

$$\nabla Q = -\frac{2}{n} (X' y - X' X \beta) \quad y = (y_1, \dots, y_n)'$$

setting equal to zero and solving to get

$$\hat{\beta} = (X' X)^{-1} X' y$$



$$\beta = \begin{pmatrix} \beta_0 \\ \vdots \\ \beta_q \end{pmatrix} \quad \nabla Q = \begin{pmatrix} \frac{\partial Q}{\partial \beta_j} \end{pmatrix}$$

2. Linear Regression and Neural Nets

$$S_i = \beta_0 + \beta_1x_{1i} + \cdots + \beta_qx_{qi}$$

A *Neural Net* is a way to perform *Multiple Linear Regression* by using a linear activation function with the corresponding normal likelihood and least squares score function.

Linear Activation

$$f(S_i) = \sum_j \beta_j x_{ji}$$

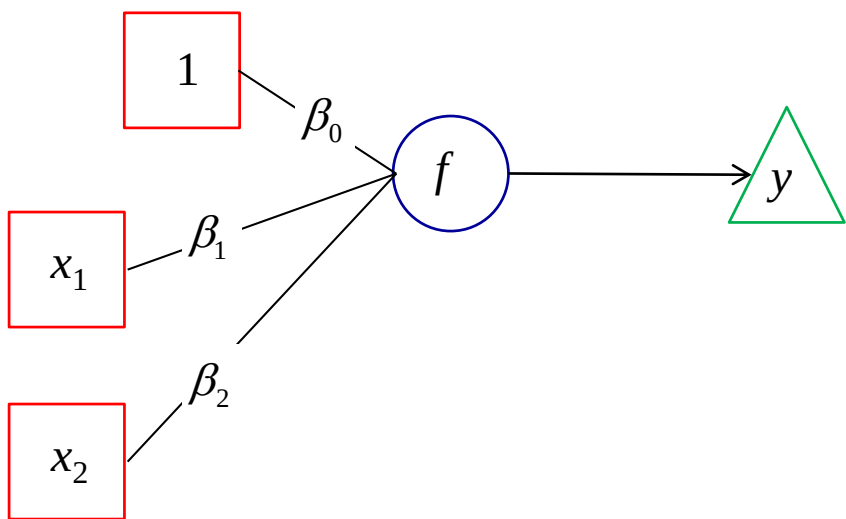
Normal Likelihood Score

$$Q = \frac{1}{n} \sum_i [y_i - \sum_j \beta_j x_{ji}]^2$$

Derivatives

$$\frac{\partial Q}{\partial \beta_j} = \frac{2}{n} \sum_i [y_i - \sum_j \beta_j x_{ij}] (-x_{ij})$$

$j = 0, 1, 2$



can set to 0 and get
 $\hat{\beta} = (X'X)^{-1}X'y$
 $y = (y_1, \dots, y_n)'$

Gradient

$$\nabla Q = -\frac{2}{n} (X'y - X'X\beta)$$

Gradient Descent

$$\hat{\beta}^{(t+1)} = \hat{\beta}^{(t)} - \gamma \nabla Q(\hat{\beta}^{(t)})$$

$$\nabla Q = \begin{bmatrix} \frac{\partial Q}{\partial \beta_0} \\ \frac{\partial Q}{\partial \beta_1} \\ \frac{\partial Q}{\partial \beta_2} \end{bmatrix}$$

Linear

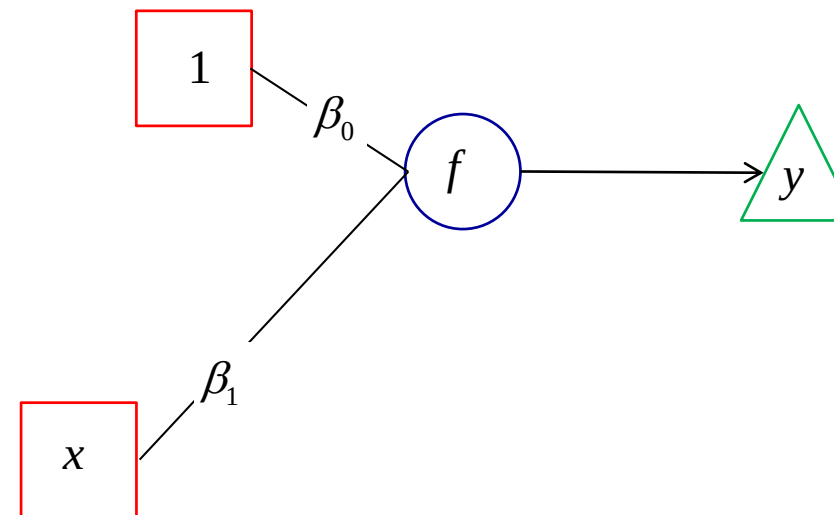
2. Linear Regression and Neural Nets

Example: Given observed data:

Data
(x,y)
(1,1.4)
(2,2.3)
(3,1.7)
(4,3.0)
(5,3.4)

Linear Activation

$$f(x) = \beta_0 + \beta_1 x$$



use the *Neural Net* structure

and *Gradient Descent* to iteratively estimate the parameters.

Normal Likelihood Score

$$Q = \frac{1}{n} \sum_i [y_i - \sum_j \beta_j x_{ji}]^2$$

Gradient

$$\nabla Q = -\frac{2}{n} (X' y - X' X \beta)$$

Gradient Descent

$$\hat{\beta}^{(t+1)} = \hat{\beta}^{(t)} - \gamma \nabla Q(\hat{\beta}^{(t)})$$

$\gamma = .0001$

2. Linear Regression and Neural Nets

Example: Given observed data:

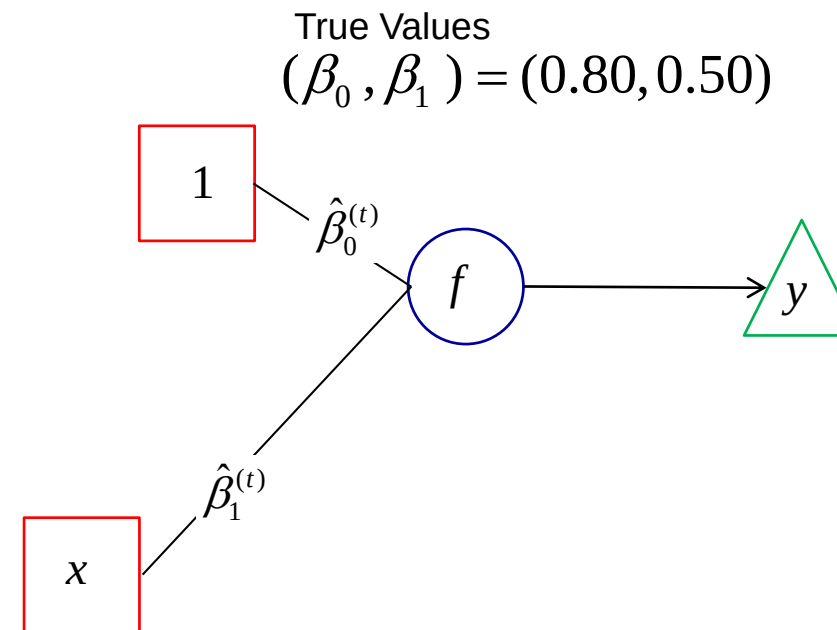
Data
(x,y)
(1,1.4)
(2,2.3)
(3,1.7)
(4,3.0)
(5,3.4)

Linear Activation

$$f(x) = \beta_0 + \beta_1 x$$

$t=0$

$$(\hat{\beta}_0^{(0)}, \hat{\beta}_1^{(0)}) = (1.50, 1.00)$$



→ Run data through with $\hat{\beta}^{(t)} = (\hat{\beta}_0^{(t)}, \hat{\beta}_1^{(t)})'$
 Calculate $\nabla Q(\hat{\beta}^{(t)}) = -\frac{2}{n}(X' y - X' X \hat{\beta}^{(t)})$, $\gamma = .0001$
 Calculate new $\hat{\beta}^{(t+1)} = \hat{\beta}^{(t)} - \gamma \nabla Q(\hat{\beta}^{(t)})$, $t = t + 1$

2. Linear Regression and Neural Nets

Example: Given observed data

Data

(x,y)
(1,1.4)
(2,2.3)
(3,1.7)
(4,3.0)
(5,3.4)

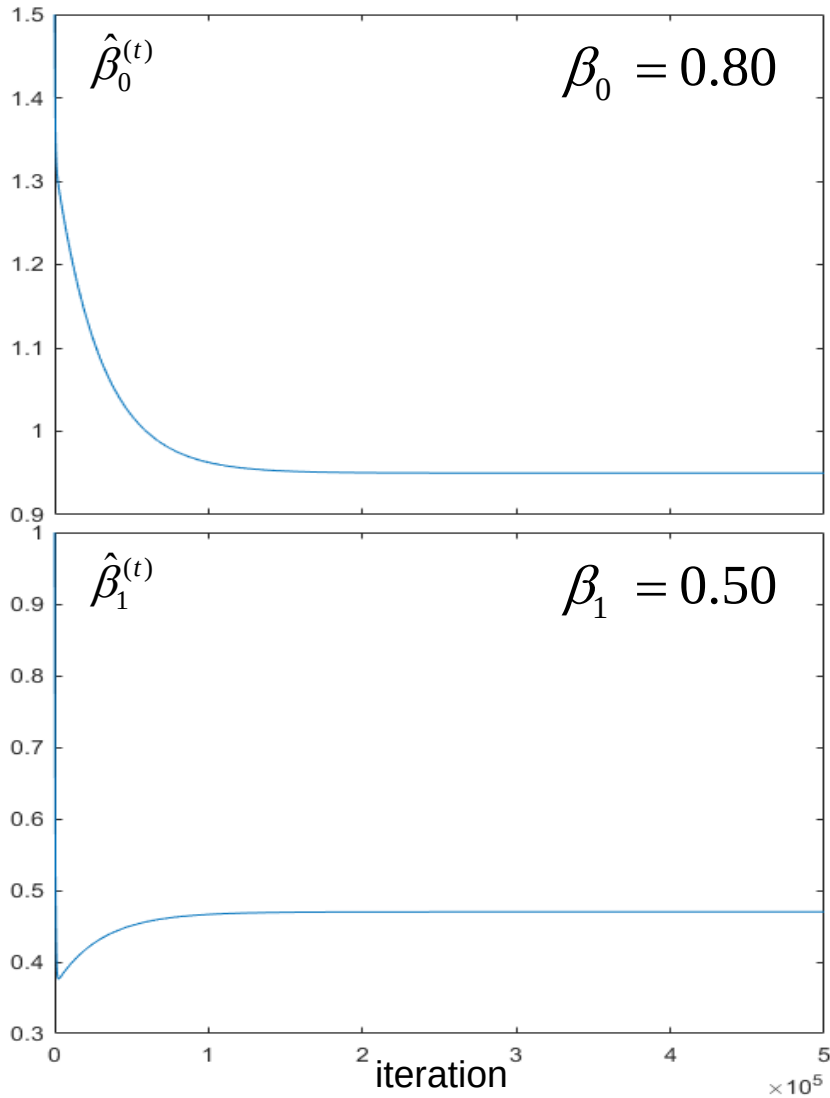
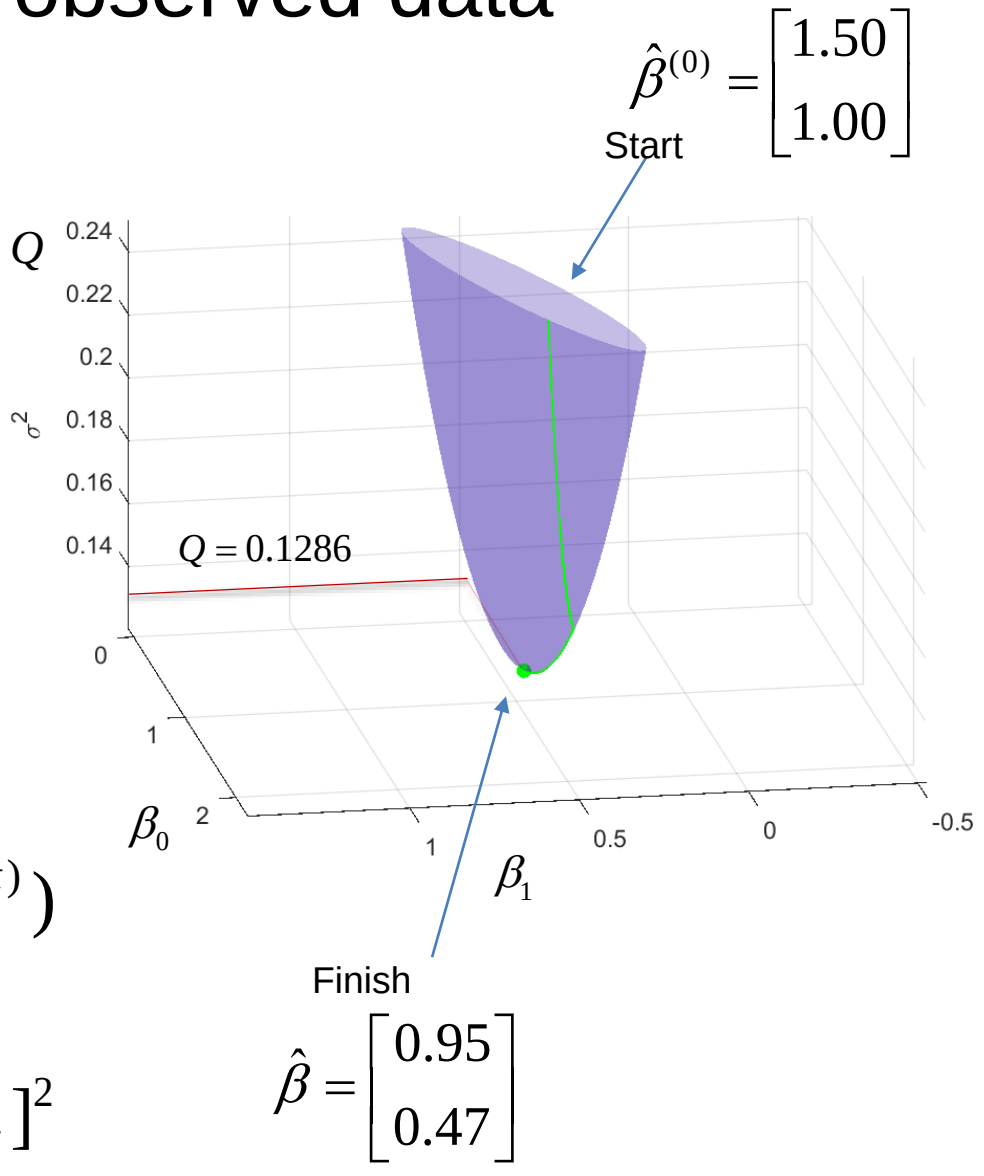
Linear Activation

$$f(x) = \beta_0 + \beta_1 x$$

$$\hat{\beta}^{(t+1)} = \hat{\beta}^{(t)} - \gamma \nabla Q(\hat{\beta}^{(t)})$$

Normal Likelihood Score

$$Q = \frac{1}{n} \sum_i [y_i - \sum_j \beta_j x_{ji}]^2$$



2. Linear Regression and Neural Nets

$$S_{ik} = \beta_{0k} + \beta_{1k}x_{1i} + \dots + \beta_{qk}x_{qi}$$

A Neural Net is a way to do *Multivariate Linear Regression* with linear activation function and normal likelihood – least squares score function.

Linear Activation

$$f_k(S_{ik}) = \sum_j \beta_{jk}x_{ji}$$

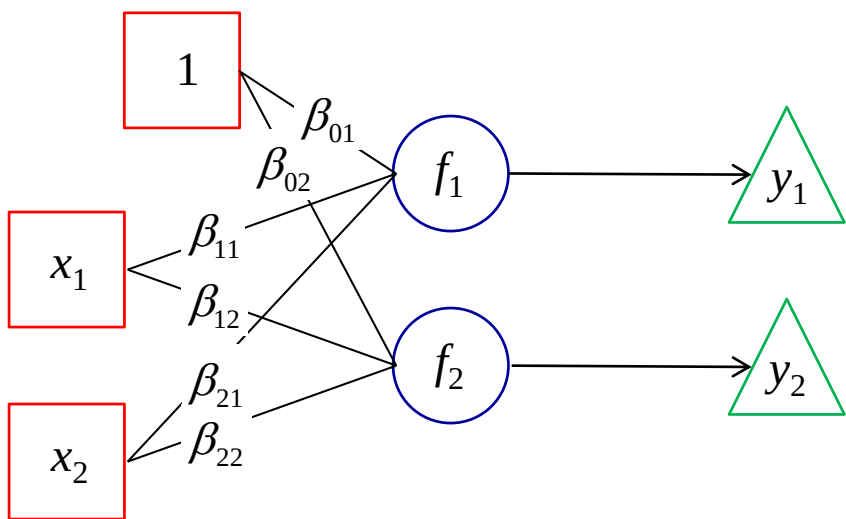
Normal Likelihood Score

$$Q_k = \frac{1}{n} \sum_i [y_i - \sum_j \beta_{jk}x_{ji}]^2$$

Derivatives

$$\frac{\partial Q_k}{\partial \beta_{jk}} = \frac{2}{n} \sum_i [y_i - \sum_j \beta_{jk}x_{ij}](-x_{ij})$$

$j = 0, 1, 2 \quad k = 1, 2$



$$\nabla Q_k = \begin{bmatrix} \frac{\partial Q_k}{\partial \beta_{0k}} \\ \frac{\partial Q_k}{\partial \beta_{1k}} \\ \frac{\partial Q_k}{\partial \beta_{2k}} \end{bmatrix}$$

Gradient

$$\nabla Q_k = -\frac{2}{n} (X' y - X' X \beta_k)$$

Gradient Descent

$$\hat{\beta}_k^{(t+1)} = \hat{\beta}_k^{(t)} - \gamma \nabla Q_k(\hat{\beta}_k^{(t)})$$

can set to 0 and get
 $\hat{\beta}_k = (X' X)^{-1} X' y_k$
 $y_k = (y_{1k}, \dots, y_{nk})'$

2. Linear Regression and Neural Nets

$$S_{i1} = \beta_{01} + \beta_{11}x_{1i} + \dots + \beta_{q1}x_{qi}$$

A Neural Net is a way to do *Multivariate Linear Regression* with linear activation function and normal likelihood – least squares score function.

Linear Activation

$$f_1(S_{i1}) = \sum_j \beta_{j1}x_{ji}$$

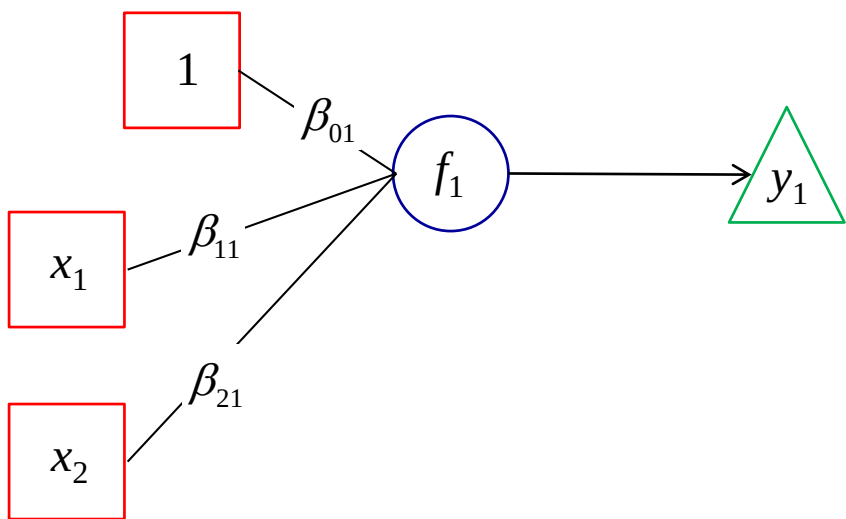
Normal Likelihood Score

$$Q_1 = \frac{1}{n} \sum_i [y_i - \sum_j \beta_{j1}x_{ji}]^2$$

Derivatives

$$\frac{\partial Q_1}{\partial \beta_{j1}} = \frac{2}{n} \sum_i [y_i - \sum_j \beta_{j1}x_{ji}](-x_{ij})$$

$j = 0, 1, 2$



$$\nabla Q_1 = \begin{bmatrix} \frac{\partial Q_1}{\partial \beta_{01}} \\ \frac{\partial Q_1}{\partial \beta_{11}} \\ \frac{\partial Q_1}{\partial \beta_{21}} \end{bmatrix}$$

Gradient

$$\nabla Q_1 = -\frac{2}{n} (X' y - X' X \beta_1)$$

Gradient Descent

$$\hat{\beta}_1^{(t+1)} = \hat{\beta}_1^{(t)} - \gamma \nabla Q_1(\hat{\beta}_1^{(t)})$$

can set to 0 and get
 $\hat{\beta}_1 = (X' X)^{-1} X' y_1$
 $y_1 = (y_{11}, \dots, y_{n1})'$

2. Linear Regression and Neural Nets

$$S_{i2} = \beta_{02} + \beta_{12}x_{1i} + \dots + \beta_{q2}x_{qi}$$

A Neural Net is a way to do *Multivariate Linear Regression* with linear activation function and normal likelihood – least squares score function.

Linear Activation

$$f_2(S_{i2}) = \sum_j \beta_{j2}x_{ji}$$

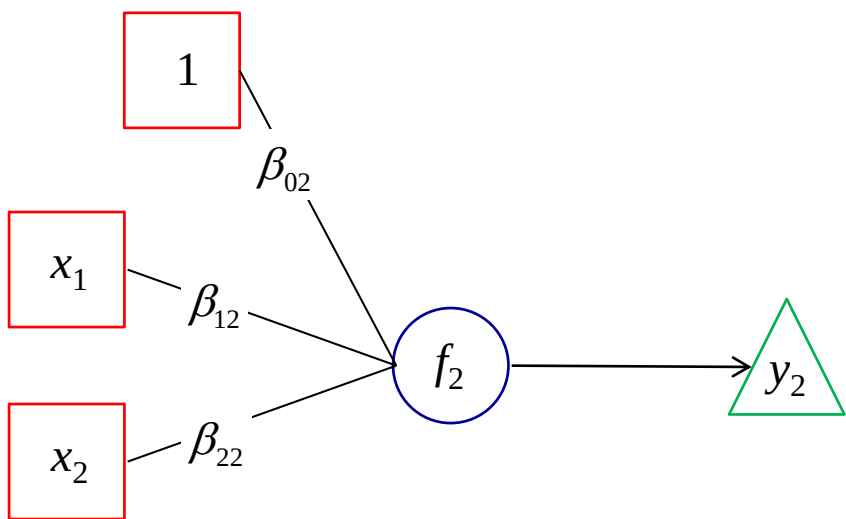
Normal Likelihood Score

$$Q_2 = \frac{1}{n} \sum_i [y_i - \sum_j \beta_{j2}x_{ji}]^2$$

Derivatives

$$\frac{\partial Q_2}{\partial \beta_{j2}} = \frac{2}{n} \sum_i [y_i - \sum_j \beta_{j2}x_{ij}](-x_{ij})$$

$j = 0, 1, 2$



$$\nabla Q_2 = \begin{bmatrix} \frac{\partial Q_2}{\partial \beta_{02}} \\ \frac{\partial Q_2}{\partial \beta_{12}} \\ \frac{\partial Q_2}{\partial \beta_{22}} \end{bmatrix}$$

Gradient

$$\nabla Q_2 = -\frac{2}{n} (X' y - X' X \beta_2)$$

Gradient Descent

$$\hat{\beta}_2^{(t+1)} = \hat{\beta}_2^{(t)} - \gamma \nabla Q_2(\hat{\beta}_2^{(t)})$$

can set to 0 and get
 $\hat{\beta}_2 = (X' X)^{-1} X' y_2$
 $y_2 = (y_{12}, \dots, y_{n2})'$

2. Linear Regression and Neural Nets

$$S_{ik} = \beta_{0k} + \beta_{1k}x_{1i} + \cdots + \beta_{qk}x_{qi}$$

Two independent parallel *Multiple Regressions* yield the same coefficient estimate as simultaneous multivariate regression (different covariance).

$$\hat{B} = (X'X)^{-1}X'Y$$

$$\hat{B} = (\hat{\beta}_1, \hat{\beta}_2)$$

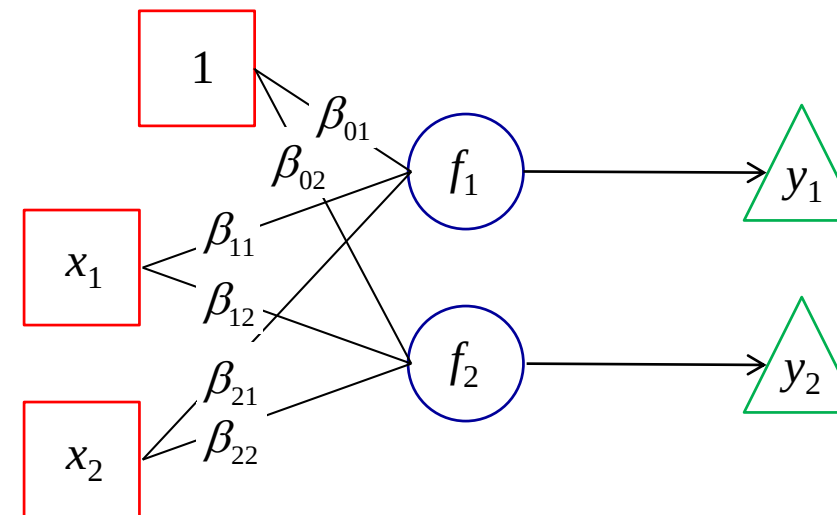
$$Y = (y_1, y_2)$$

$$\hat{\beta}_1 = (X'X)^{-1}X'y_1$$

$$\hat{\beta}_2 = (X'X)^{-1}X'y_2$$

$$y_1 = (y_{11}, \dots, y_{n1})'$$

$$y_2 = (y_{12}, \dots, y_{n2})'$$



$$\hat{\beta}_1^{(t+1)} = \hat{\beta}_1^{(t)} - \gamma \nabla Q_1(\hat{\beta}_1^{(t)})$$

$$\hat{\beta}_2^{(t+1)} = \hat{\beta}_2^{(t)} - \gamma \nabla Q_2(\hat{\beta}_2^{(t)})$$

Now the trained *Neural Net* can be applied to new data.

3. Non-Linear Regression and Neural Nets

We might believe that there is a non-linear relationship between an independent variable x , and a dependent variable y with measurement error.

$$y_i = f(x_i) + \varepsilon_i \quad i = 1, \dots, n$$

We can assume normal errors or use least squares.

$$Q = \frac{1}{n} \sum_{i=1}^n (y_i - f(x_i))^2$$

As an example consider the SoftPlus function $f(\cdot) = \ln(1 + e^{\sum_j \beta_j x_j})$.

3. Non-Linear Regression and Neural Nets

$$S_i = \beta_0 + \beta_1x_{1i} + \cdots + \beta_qx_{qi}$$

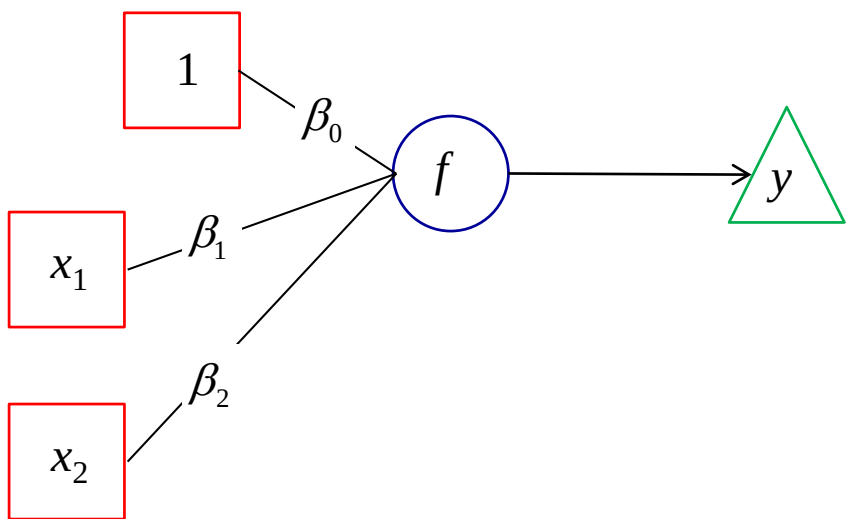
A Neural Net is a way to do *Non-Linear Regression* with SoftPlus activation function and normal likelihood – least squares score function.

SoftPlus Activation

$$f(S_i) = \ln(1 + e^{S_i})$$

Normal Likelihood Score

$$Q = \frac{1}{n} \sum_i [y_i - \ln(1 + e^{S_i})]^2$$



$$\nabla Q = \begin{bmatrix} \frac{\partial Q}{\partial \beta_0} \\ \frac{\partial Q}{\partial \beta_1} \\ \frac{\partial Q}{\partial \beta_2} \end{bmatrix}$$

Derivatives

$$\frac{\partial Q}{\partial \beta_j} = \frac{2}{n} \sum_i \left[\frac{x_{ij} [y_i - \ln(1 + \exp(\sum_j \beta_j x_{ij}))] \exp(\sum_j \beta_j x_{ij})}{1 + \exp(\sum_j \beta_j x_{ij})} \right]$$

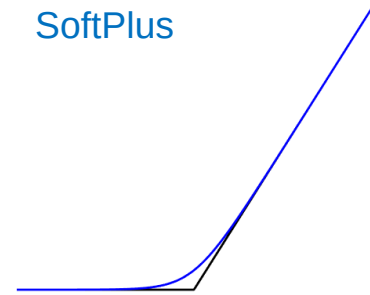
Gradient

$$\nabla Q = \left(\frac{\partial Q}{\partial \beta_j} \right)$$

Gradient Descent

$$\hat{\beta}^{(t+1)} = \hat{\beta}^{(t)} - \gamma \nabla Q(\hat{\beta}^{(t)})$$

SoftPlus



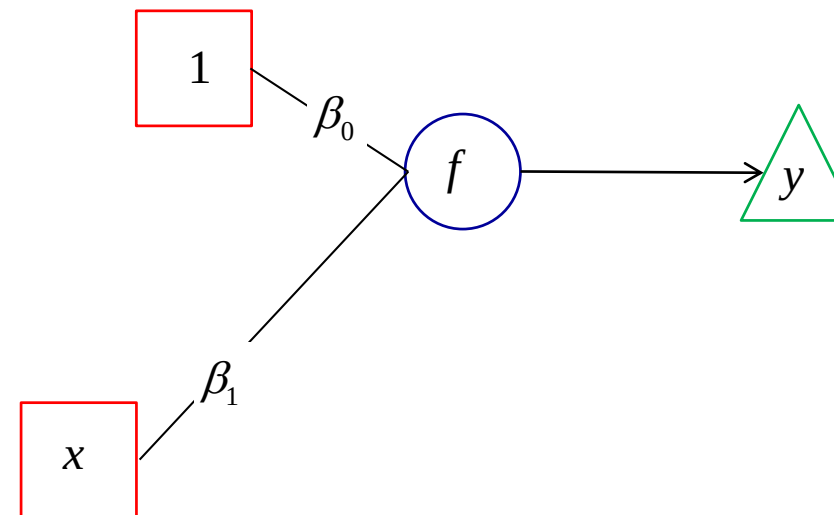
3. Non-Linear Regression and Neural Nets

Example: Given observed data:

$$\frac{\partial Q}{\partial \beta_j} = \frac{2}{n} \sum_i \left[\frac{x_{ij} [y_i - \ln(1 + \exp(\sum_j \beta_j x_{ij}))] \exp(\sum_j \beta_j x_{ij})}{1 + \exp(\sum_j \beta_j x_{ij})} \right]$$

SoftPlus Activation

$$f(S_i) = \ln(1 + e^{\beta_0 + \beta_1 x_{1i}})$$



use the *Neural Net* structure

and *Gradient Descent* to iteratively estimate the parameters.

Normal Likelihood Score

$$Q = \frac{1}{n} \sum_i [y_i - \ln(1 + e^{\beta_0 + \beta_1 x_{1i}})]^2$$

Gradient

$$\nabla Q = \left(\frac{\partial Q}{\partial \beta_j} \right)$$

Gradient Descent

$$\hat{\beta}^{(t+1)} = \hat{\beta}^{(t)} - \gamma \nabla Q(\hat{\beta}^{(t)})$$

$\gamma = .0001$

3. Non-Linear Regression and Neural Nets

Example: Given observed data:

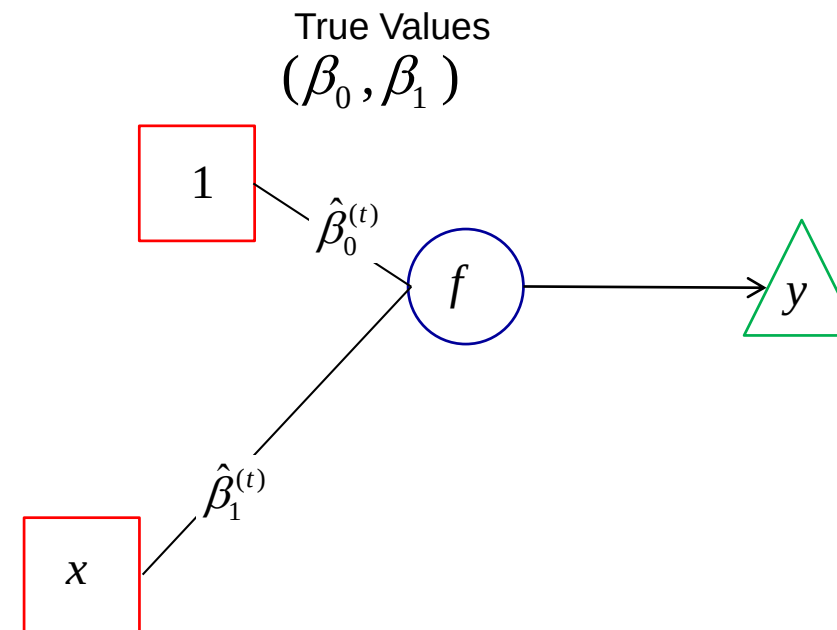
$$\frac{\partial Q}{\partial \beta_j} = \frac{2}{n} \sum_i \left[\frac{x_{ij} [y_i - \ln(1 + \exp(\sum_j \beta_j x_{ij}))] \exp(\sum_j \beta_j x_{ij})}{1 + \exp(\sum_j \beta_j x_{ij})} \right]$$

Linear Activation

$$f(x) = \beta_0 + \beta_1 x$$

$t=0$

$$(\hat{\beta}_0^{(0)}, \hat{\beta}_1^{(0)})$$



→ Run data through with $\hat{\beta}^{(t)} = (\hat{\beta}_0^{(t)}, \hat{\beta}_1^{(t)})'$

Calculate $\nabla Q(\hat{\beta}^{(t)})$, $\gamma = .0001$

Calculate new $\hat{\beta}^{(t+1)} = \hat{\beta}^{(t)} - \gamma \nabla Q(\hat{\beta}^{(t)})$, $t=t+1$ →

3. Non-Linear Regression and Neural Nets

$$S_{ik} = \beta_{0k} + \beta_{1k}x_{1i} + \dots + \beta_{qk}x_{qi}$$

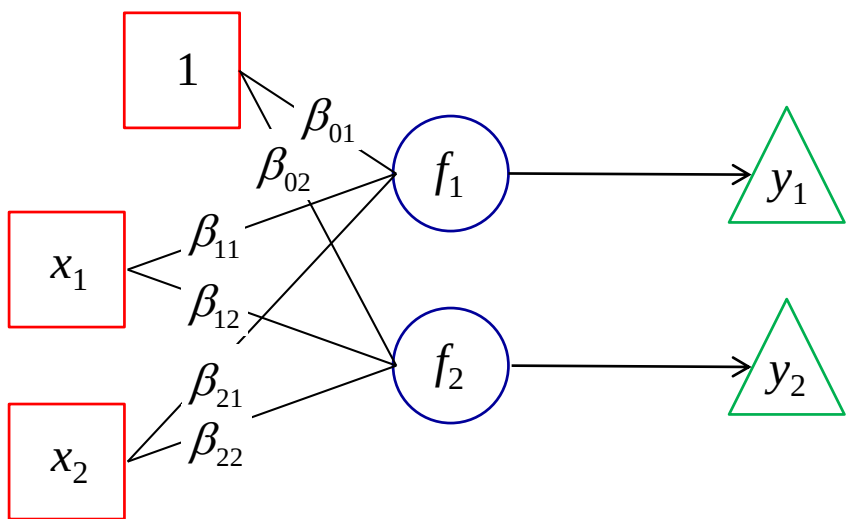
A *Neural Net* is a way to do *Multivariate Non-Linear Regression* with SoftPlus activation function and normal likelihood – least squares score function.

SoftPlus Activation

$$f(S_{ik}) = \ln(1 + e^{S_{ik}})$$

Normal Likelihood Score

$$Q_k = \frac{1}{n} \sum_i [y_i - \ln(1 + e^{S_{ik}})]^2$$



$$\nabla Q_k = \begin{bmatrix} \frac{\partial Q_k}{\partial \beta_{0k}} \\ \frac{\partial Q_k}{\partial \beta_{1k}} \\ \frac{\partial Q_k}{\partial \beta_{2k}} \end{bmatrix}$$

Derivatives

$$\frac{\partial Q_k}{\partial \beta_{jk}} = \frac{2}{n} \sum_i \left[\frac{x_{ij} [y_{ik} - \ln(1 + \exp(\sum_j \beta_{jk} x_{ij}))] \exp(\sum_j \beta_{jk} x_{ij})}{1 + \exp(\sum_j \beta_{jk} x_{ij})} \right]_{k=1,2}$$

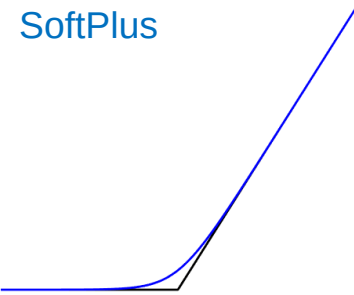
Gradient

$$\nabla Q_k = \left(\frac{\partial Q_k}{\partial \beta_{jk}} \right)$$

Gradient Descent

$$\hat{\beta}_k^{(t+1)} = \hat{\beta}_k^{(t)} - \gamma \nabla Q_k(\hat{\beta}_k^{(t)})$$

SoftPlus



3. Non-Linear Regression and Neural Nets

$$S_{ik} = \beta_{0k} + \beta_{1k}x_{1i} + \dots + \beta_{qk}x_{qi}$$

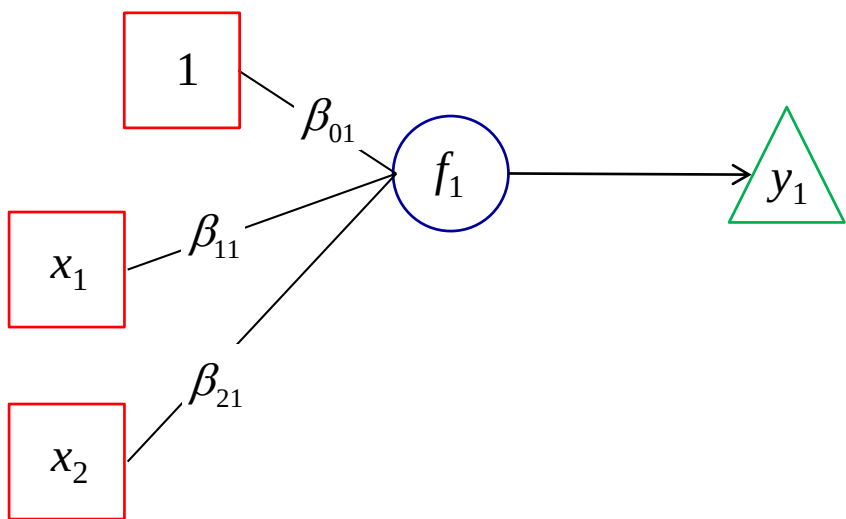
A *Neural Net* is a way to do *Multivariate Non-Linear Regression* with SoftPlus activation function and normal likelihood – least squares score function.

SoftPlus Activation

$$f(S_{i1}) = \ln(1 + e^{S_{i1}})$$

Normal Likelihood Score

$$Q_1 = \frac{1}{n} \sum_i [y_i - \ln(1 + e^{S_{i1}})]^2$$



$$\nabla Q_1 = \begin{bmatrix} \frac{\partial Q_1}{\partial \beta_{01}} \\ \frac{\partial Q_1}{\partial \beta_{11}} \\ \frac{\partial Q_1}{\partial \beta_{21}} \end{bmatrix}$$

Derivatives

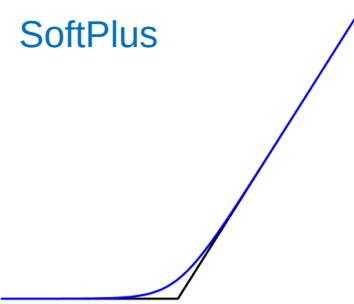
$$\frac{\partial Q_1}{\partial \beta_{j1}} = \frac{2}{n} \sum_i \left[\frac{x_{ij} [y_{i1} - \ln(1 + \exp(\sum_j \beta_{j1} x_{ij}))] \exp(\sum_j \beta_{j1} x_{ij})}{1 + \exp(\sum_j \beta_{j1} x_{ij})} \right]$$

Gradient

$$\nabla Q_1 = \left(\frac{\partial Q_1}{\partial \beta_{j1}} \right)$$

Gradient Descent

$$\hat{\beta}_1^{(t+1)} = \hat{\beta}_1^{(t)} - \gamma \nabla Q_1(\hat{\beta}_1^{(t)})$$



3. Non-Linear Regression and Neural Nets

$$S_{i2} = \beta_{02} + \beta_{12}x_{1i} + \dots + \beta_{q2}x_{qi}$$

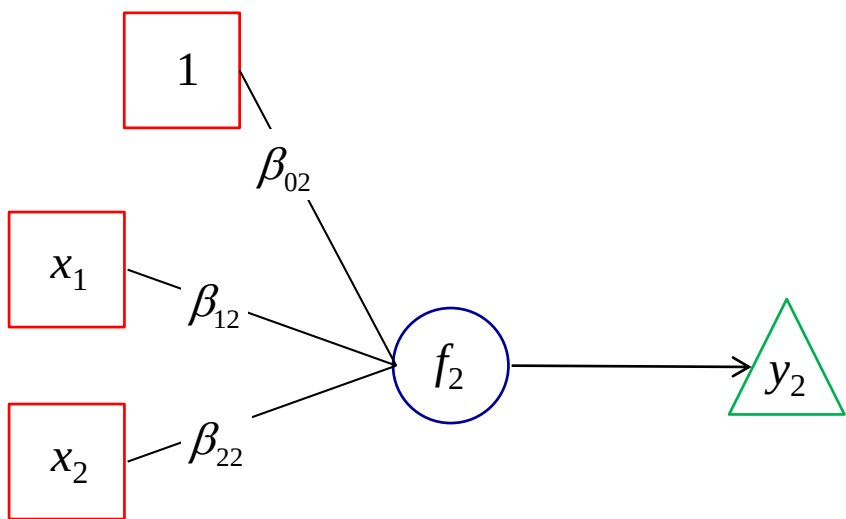
A *Neural Net* is a way to do *Multivariate Non-Linear Regression* with SoftPlus activation function and normal likelihood – least squares score function.

SoftPlus Activation

$$f(S_{i2}) = \ln(1 + e^{S_{i2}})$$

Normal Likelihood Score

$$Q_2 = \frac{1}{n} \sum_i [y_i - \ln(1 + e^{S_{i2}})]^2$$



$$\nabla Q_2 = \begin{bmatrix} \frac{\partial Q_2}{\partial \beta_{02}} \\ \frac{\partial Q_2}{\partial \beta_{12}} \\ \frac{\partial Q_2}{\partial \beta_{22}} \end{bmatrix}$$

Derivatives

$$\frac{\partial Q_2}{\partial \beta_{j2}} = \frac{2}{n} \sum_i \left[\frac{x_{ij} [y_{i2} - \ln(1 + \exp(\sum_j \beta_{j2} x_{ij}))] \exp(\sum_j \beta_{j2} x_{ij})}{1 + \exp(\sum_j \beta_{j2} x_{ij})} \right]$$

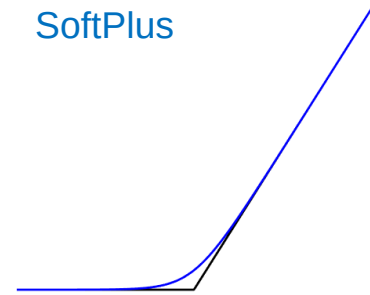
Gradient

$$\nabla Q_2 = \left(\frac{\partial Q_2}{\partial \beta_{j2}} \right)$$

Gradient Descent

$$\hat{\beta}_2^{(t+1)} = \hat{\beta}_2^{(t)} - \gamma \nabla Q_2(\hat{\beta}_2^{(t)})$$

SoftPlus



3. Non-Linear Regression and Neural Nets

$$S_{ik} = \beta_{0k} + \beta_{1k}x_{1i} + \dots + \beta_{qk}x_{qi}$$

Two independent Non-Linear Regressions yield the same coefficient estimate as simultaneous Non-Linear Regression.

$$\hat{B} = (\hat{\beta}_1, \hat{\beta}_2)$$

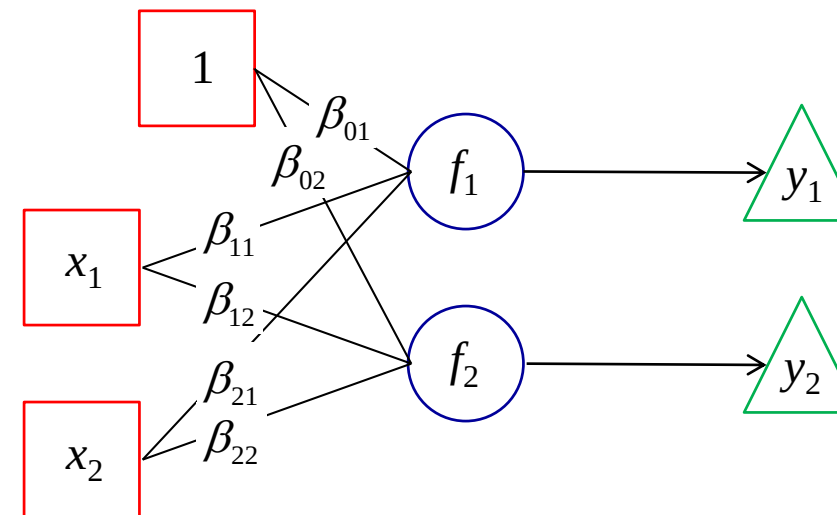
$$Y = (y_1, y_2)$$

$\hat{\beta}_1$ The last one in iteration.

$\hat{\beta}_2$ The last one in iteration.

$$y_1 = (y_{11}, \dots, y_{n1})'$$

$$y_2 = (y_{12}, \dots, y_{n2})'$$



$$\hat{\beta}_1^{(t+1)} = \hat{\beta}_1^{(t)} - \gamma \nabla Q_1(\hat{\beta}_1^{(t)})$$

$$\hat{\beta}_2^{(t+1)} = \hat{\beta}_2^{(t)} - \gamma \nabla Q_2(\hat{\beta}_2^{(t)})$$

Now the trained *Neural Net* can be applied to new data.

4. Logistic Regression and Neural Nets

Often the probability p of an event E depends upon an independent variable x , such as the probability of getting an A on a class final depends on the number of hours that a student studies x .

So p is a function of x , $p(x)$. $0 \leq p(x) \leq 1$.

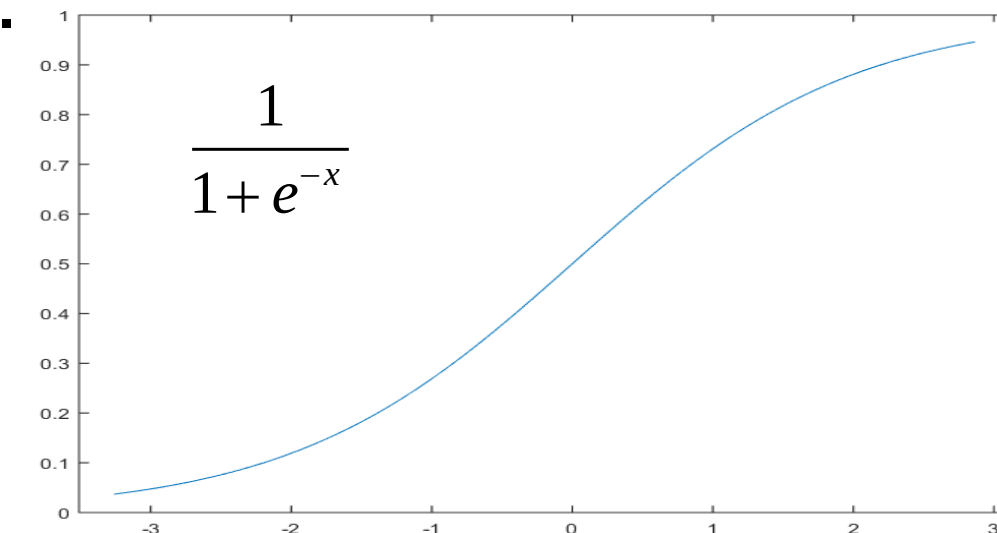
Hours (x)	A (y)
6	0
8	0
10	0
12	0
14	0
16	1
18	0
20	0
22	0
24	0
26	1
28	0
30	0
32	1
34	1
36	1
38	1
40	1

4. Logistic Regression and Neural Nets

This dependency of a probability $p(x)$, $0 \leq p(x) \leq 1$, on an independent variable x , $-\infty < x < \infty$, is generally described through a *link function*, here the logistic mapping function

$$p = p(x) = \frac{1}{1 + e^{-(\beta_0 + \beta_1 x)}} .$$

If the event E occurs, then we say $y=1$ and if not, $y=0$. $P(y=1)=p$ and $P(y=0)=1-p$ This is a Bernoulli trial.



4. Logistic Regression and Neural Nets

The likelihood function

$$L(\beta_0, \beta_1) = \prod_{i=1}^n [p(x_i)]^{y_i} [1 - p(x_i)]^{1-y_i} \quad \begin{array}{l} y_i = \{0,1\} \\ -\infty < x_i < \infty \end{array}$$

where $p(x_i) = \frac{1}{1 + e^{-\beta_0 - \beta_1 x_i}}$

has log likelihood function

$$\begin{aligned} LL(\beta_0, \beta_1) &= \sum_{i=1}^n [y_i \ln[p(x_i)] + (1 - y_i) \ln[1 - p(x_i)]] \\ &= \sum_{i=1}^n y_i \ln[p(x_i)] + \sum_{i=1}^n (1 - y_i) \ln[1 - p(x_i)] \\ &= \sum_{i=1}^n \ln[1 - p(x_i)] + \sum_{i=1}^n y_i \ln[p(x_i) / (1 - p(x_i))] \\ &= \sum_{i=1}^n y_i [\beta_0 + \beta_1 x_i] - \sum_{i=1}^n \ln[1 + e^{\beta_0 + \beta_1 x_i}] . \end{aligned}$$

Caution!

Do not use least squares!

~~$$Q = \frac{1}{n} \sum_i \left[y_i - \frac{1}{1 + e^{-\beta_0 - \beta_1 x_i}} \right]^2$$~~

4. Logistic Regression and Neural Nets

We can estimate the “best” logistic relationship between x and 0/1 y using the score function

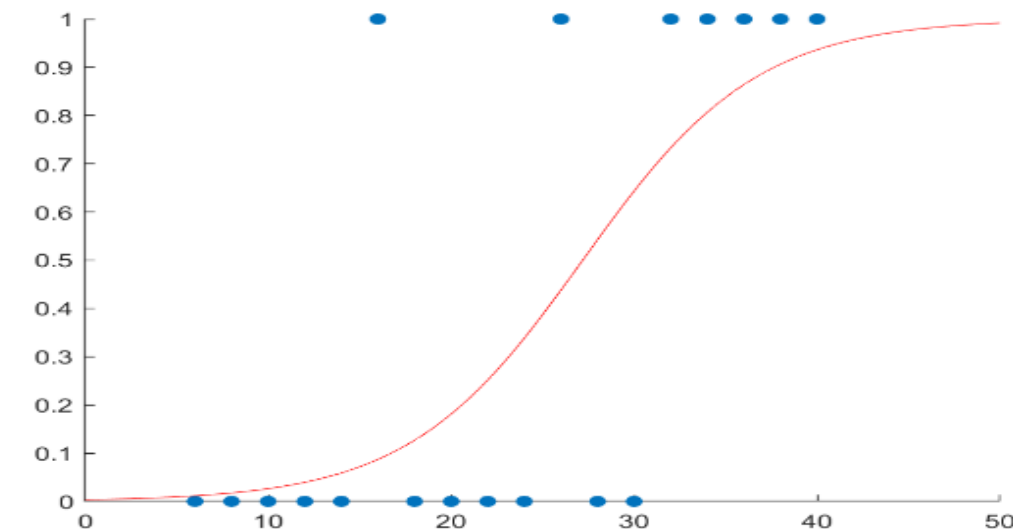
$$Q = \sum_i y_i \left(\sum_j \beta_j x_{ij} \right) - \sum_i \ln[1 + \exp(\sum_j \beta_j x_{ij})]$$

by taking derivatives with respect to the beta's

$$\frac{\partial Q}{\partial \beta_j} = \sum_i x_{ij} y_i - \sum_i \frac{x_{ij}}{1 + \exp(-\sum_j \beta_j x_{ij})} \quad \begin{array}{l} i = 1, \dots, n \quad j = 1, \dots, q \\ x_{i0} = 1 \end{array}$$

with no closed form solution.

$$f(\cdot) = \frac{1}{1 + \exp(-\sum_j \beta_j x_j)}$$



4. Logistic Regression and Neural Nets

A Neural Net is a way to represent *Multiple Logistic Regression* with Logistic activation and Bernoulli likelihood score function.

Logistic Activation

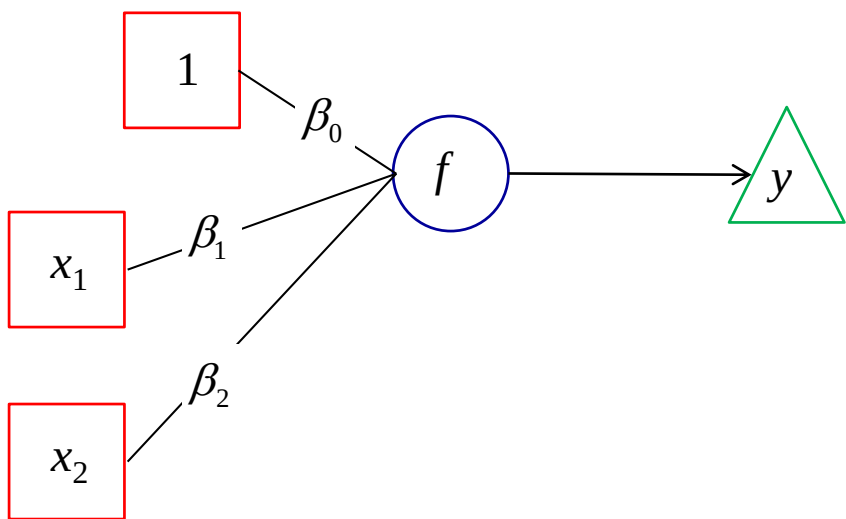
$$f(x) = \frac{1}{1 + \exp(-\sum_j \beta_j x_j)}$$

Bernoulli Likelihood Score

$$Q = \sum_i y_i (\sum_j \beta_j x_{ji}) - \sum_i \ln[1 + \exp(\sum_j \beta_j x_{ji})]$$

Derivatives

$$\frac{\partial Q}{\partial \beta_j} = \sum_i x_{ij} y_i - \sum_i \frac{x_{ij}}{1 + \exp(-\sum_j \beta_j x_{ij})}$$



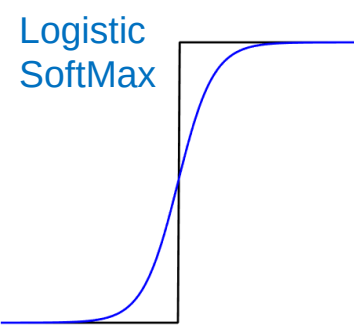
Gradient

$$\nabla Q = \left(\frac{\partial Q}{\partial \beta_j} \right)$$

Gradient Descent

$$\hat{\beta}^{(t+1)} = \hat{\beta}^{(t)} - \gamma \nabla Q(\hat{\beta}^{(t)})$$

$$\nabla Q = \begin{bmatrix} \frac{\partial Q}{\partial \beta_0} \\ \frac{\partial Q}{\partial \beta_1} \\ \frac{\partial Q}{\partial \beta_2} \end{bmatrix}$$

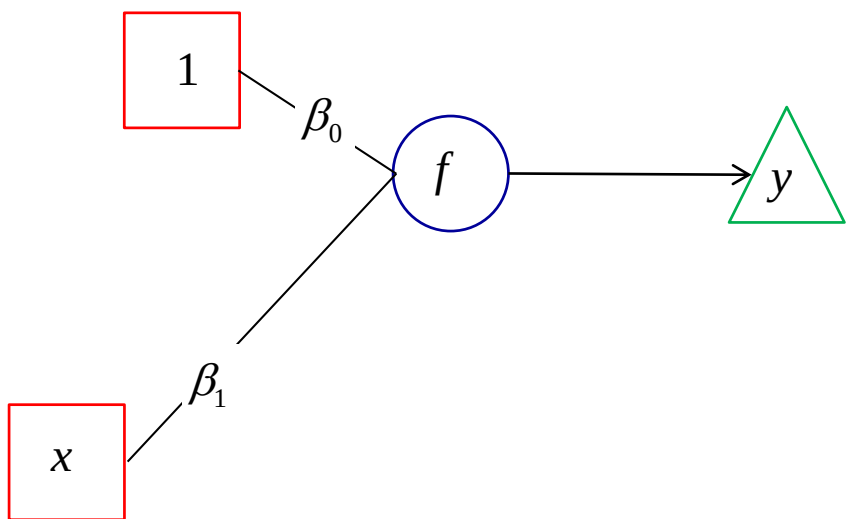


4. Logistic Regression and Neural Nets

Simple *Logistic Regression*
Given observed data:

$$f(x) = \frac{1}{1 + \exp(-\sum_j \beta_j x_j)}$$

use the *Neural Net* structure
and *Gradient Descent* to
iteratively estimate the parameters.



Gradient

$$\nabla Q = \left(\frac{\partial Q}{\partial \beta_j} \right)$$

Gradient Descent

$$\hat{\beta}^{(t+1)} = \hat{\beta}^{(t)} - \gamma \nabla Q(\hat{\beta}^{(t)})$$

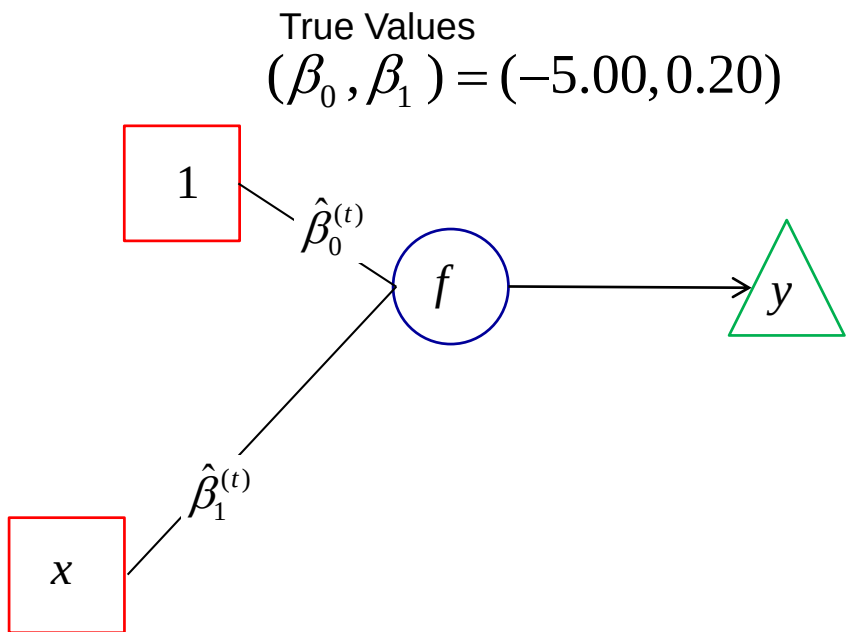
Hours (x)	A (y)
6	0
8	0
10	0
12	0
14	0
16	1
18	0
20	0
22	0
24	0
26	1
28	0
30	0
32	1
34	1
36	1
38	1
40	1

4. Logistic Regression and Neural Nets

Simple *Logistic Regression*
Given observed data:

$$f(x) = \frac{1}{1 + \exp(-\sum_j \beta_j x_j)}$$

$t=0$
 $(\hat{\beta}_0^{(0)}, \hat{\beta}_1^{(0)}) = (3.00, 0.50)$



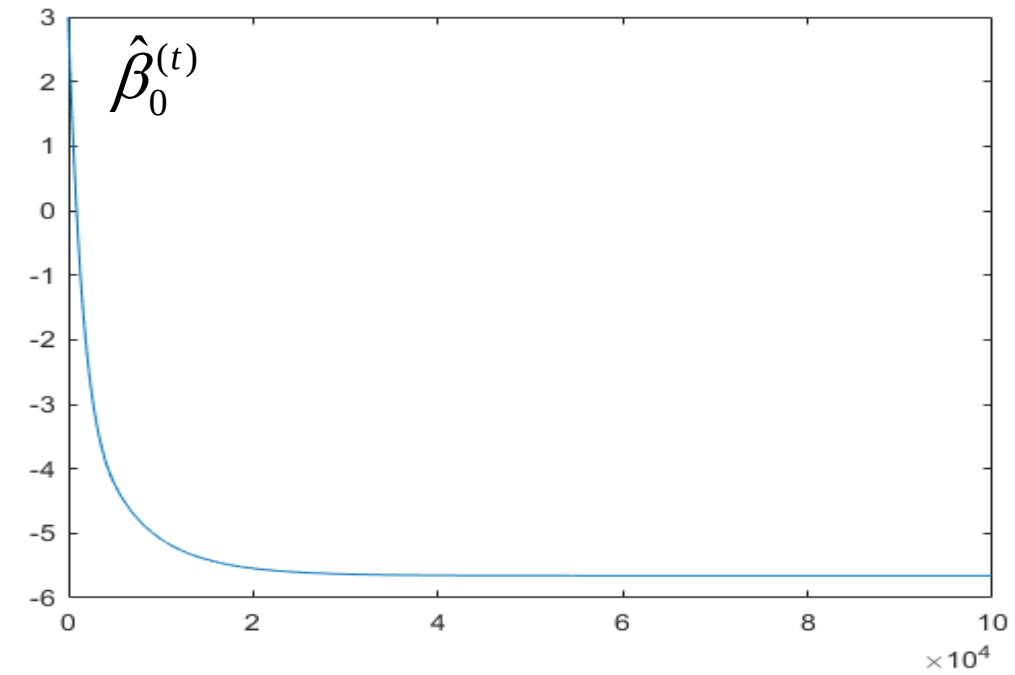
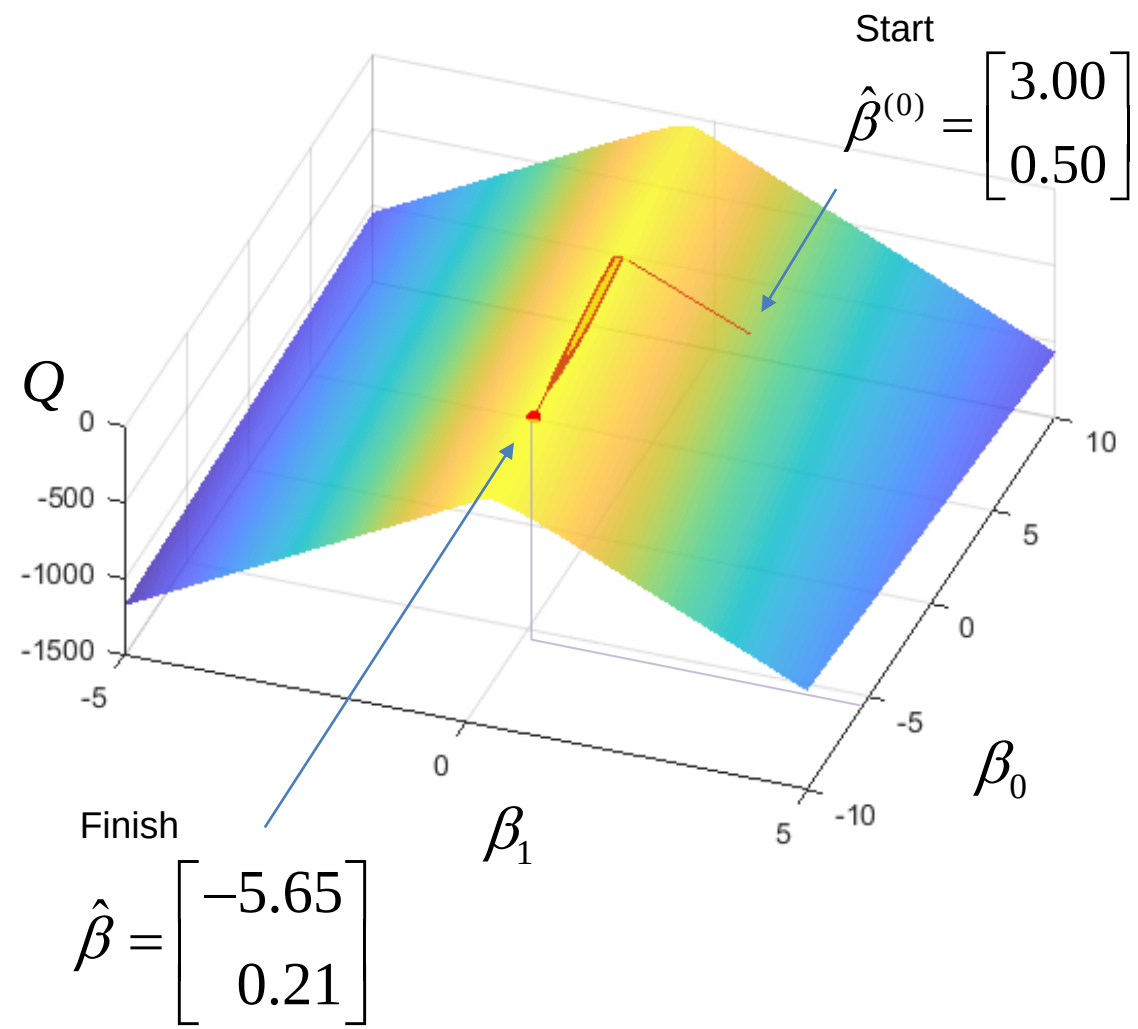
Hours (x)	A (y)
6	0
8	0
10	0
12	0
14	0
16	1
18	0
20	0
22	0
24	0
26	1
28	0
30	0
32	1
34	1
36	1
38	1
40	1

→ Run data through with $\hat{\beta}^{(t)} = (\hat{\beta}_0^{(t)}, \hat{\beta}_1^{(t)})'$
Calculate $\nabla Q(\hat{\beta}^{(t)}) = \sum_i x_{ij} y_i - \sum_i \frac{x_{ij}}{1 + \exp(-\sum_j \beta_j x_{ij})}$, $\gamma = .0001$
Calculate new $\hat{\beta}^{(t+1)} = \hat{\beta}^{(t)} - \gamma \nabla Q(\hat{\beta}^{(t)})$, $t \pm t+1$

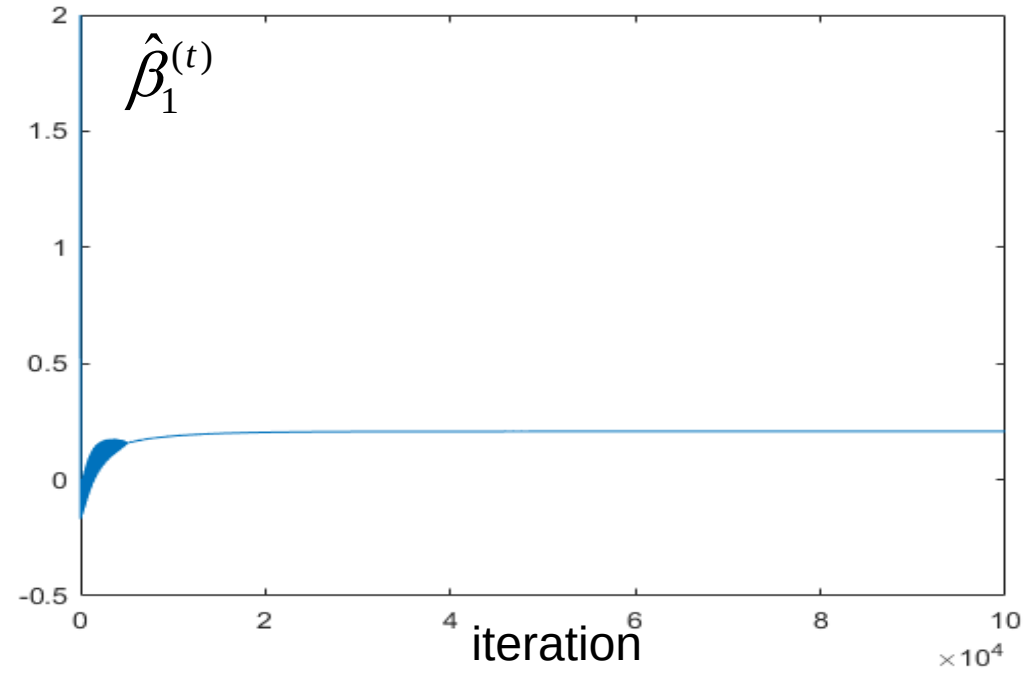
4. Logistic Regression and Neural Nets

Simple *Logistic Regression* results:

$\beta_0 = -5.00$



$\beta_1 = 0.20$



4. Logistic Regression and Neural Nets

A Neural Net is a way to represent *Multivariate Logistic Regression* with logistic activation and Bernoulli score function.

Logistic Activation

$$f_k(x) = \frac{1}{1 + \exp(-\sum_j \beta_{jk} x_j)}$$

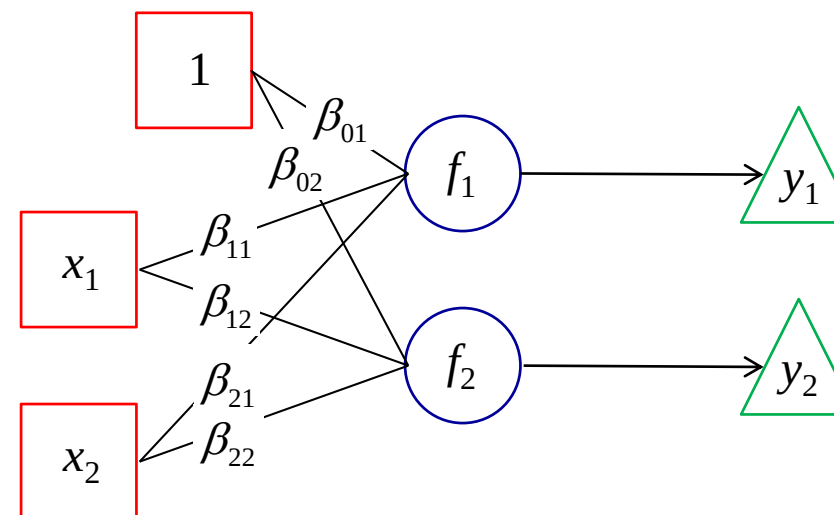
Bernoulli Likelihood Score

$$Q_k = \sum_i y_{ik} (\sum_j \beta_{jk} x_{ij}) - \sum_i \ln[1 + \exp(\sum_j \beta_{jk} x_{ij})]$$

Derivatives

$$\frac{\partial Q_k}{\partial \beta_{jk}} = \sum_i x_{ij} y_{ik} - \sum_i \frac{x_{ij}}{1 + \exp(-\sum_j \beta_{jk} x_{ij})}$$

$j = 0, 1, 2 \quad k = 1, 2$



Gradient

$$\nabla Q_k = \left(\frac{\partial Q_k}{\partial \beta_{jk}} \right)$$

Gradient Descent

$$\hat{\beta}_k^{(t+1)} = \hat{\beta}_k^{(t)} - \gamma \nabla Q_k(\hat{\beta}_k^{(t)})$$

$$\nabla Q_k = \begin{bmatrix} \frac{\partial Q_k}{\partial \beta_{0k}} \\ \frac{\partial Q_k}{\partial \beta_{1k}} \\ \frac{\partial Q_k}{\partial \beta_{2k}} \end{bmatrix}$$

4. Logistic Regression and Neural Nets

A Neural Net is a way to represent *Multivariate Logistic Regression* with logistic activation and Bernoulli score function.

Logistic Activation

$$f_1(x) = \frac{1}{1 + \exp(-\sum_j \beta_{j1} x_j)}$$

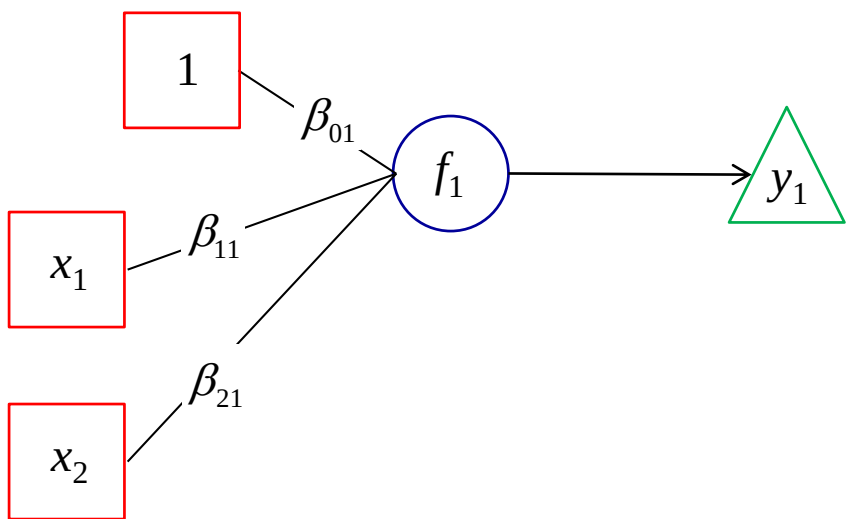
Bernoulli Likelihood Score

$$Q_1 = \sum_i y_{i1} (\sum_j \beta_{j1} x_{ij}) - \sum_i \ln[1 + \exp(\sum_j \beta_{j1} x_{ij})]$$

Derivatives

$$\frac{\partial Q_1}{\partial \beta_{j1}} = \sum_i x_{ij} y_{i1} - \sum_i \frac{x_{ij}}{1 + \exp(-\sum_j \beta_{j1} x_{ij})}$$

$j = 0, 1, 2$



$$\nabla Q_1 = \begin{bmatrix} \frac{\partial Q_1}{\partial \beta_{01}} \\ \frac{\partial Q_1}{\partial \beta_{11}} \\ \frac{\partial Q_1}{\partial \beta_{21}} \end{bmatrix}$$

Gradient

$$\nabla Q_k = \left(\frac{\partial Q_k}{\partial \beta_{jk}} \right)$$

Gradient Descent

$$\hat{\beta}_k^{(t+1)} = \hat{\beta}_k^{(t)} - \gamma \nabla Q_k(\hat{\beta}_k^{(t)})$$

4. Logistic Regression and Neural Nets

A Neural Net is a way to represent *Multivariate Logistic Regression* with logistic activation and Bernoulli score function.

Logistic Activation

$$f_2(x) = \frac{1}{1 + \exp(-\sum_j \beta_{j2} x_j)}$$

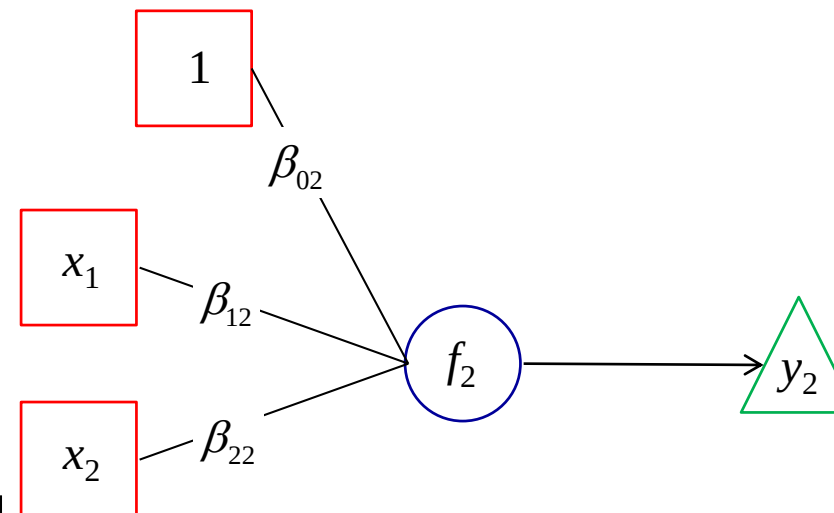
Bernoulli Likelihood Score

$$Q_2 = \sum_i y_{i2} (\sum_j \beta_{j2} x_{ij}) - \sum_i \ln[1 + \exp(\sum_j \beta_{j2} x_{ij})]$$

Derivatives

$$\frac{\partial Q_2}{\partial \beta_{j2}} = \sum_i x_{ij} y_{i2} - \sum_i \frac{x_{ij}}{1 + \exp(-\sum_j \beta_{j2} x_{ij})}$$

$j = 0, 1, 2$



Gradient

$$\nabla Q_2 = \left(\frac{\partial Q_2}{\partial \beta_{j2}} \right)$$

Gradient Descent

$$\hat{\beta}_2^{(t+1)} = \hat{\beta}_2^{(t)} - \gamma \nabla Q_2(\hat{\beta}_2^{(t)})$$

$$\nabla Q_2 = \begin{bmatrix} \frac{\partial Q_2}{\partial \beta_{02}} \\ \frac{\partial Q_2}{\partial \beta_{12}} \\ \frac{\partial Q_2}{\partial \beta_{22}} \end{bmatrix}$$

4. Logistic Regression and Neural Nets

Two independent parallel regressions yield the “*Multivariate*” *Logistic Regression* results we are after.

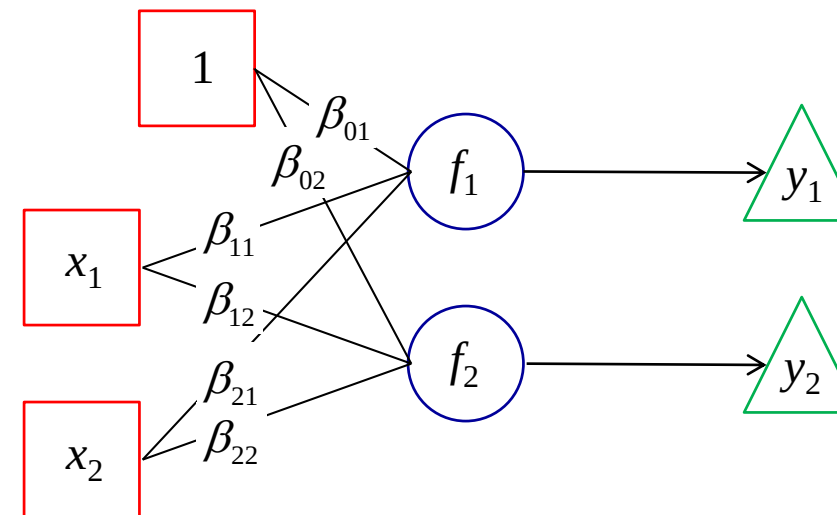
Estimated Coefficients

$$\hat{B} = (\hat{\beta}_1, \hat{\beta}_2)$$

Maximizing Score Functions

$$Q_1 = \sum_i y_{i1} \left(\sum_j \beta_{j1} x_{ij} \right) - \sum_i \ln[1 + \exp(\sum_j \beta_{j1} x_{ij})]$$

$$Q_2 = \sum_i y_{i2} \left(\sum_j \beta_{j2} x_{ij} \right) - \sum_i \ln[1 + \exp(\sum_j \beta_{j2} x_{ij})]$$



Independently
Estimated
Via Gradient
Descent

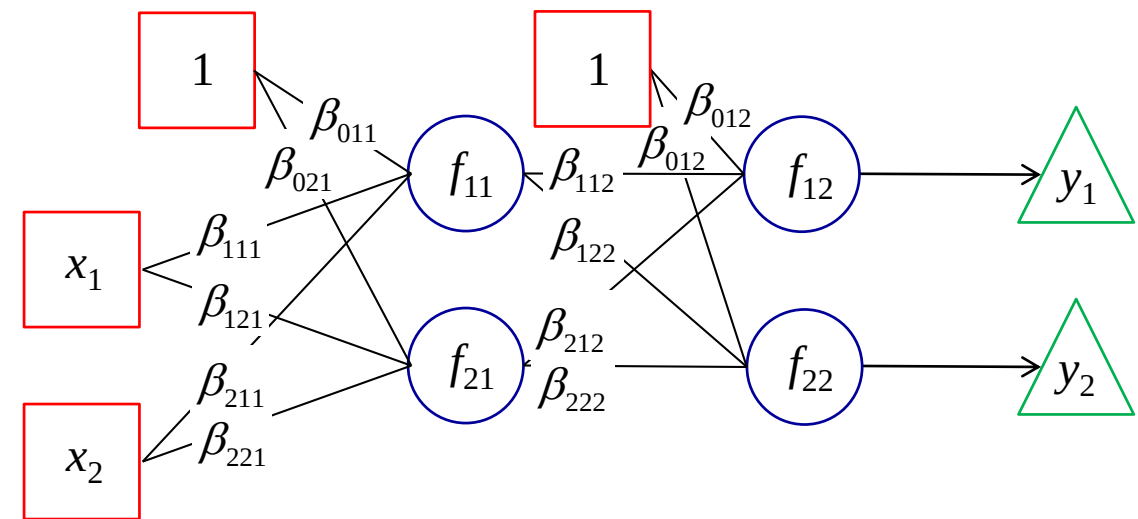
$$\hat{\beta}_1^{(t+1)} = \hat{\beta}_1^{(t)} + \gamma \nabla Q_1(\hat{\beta}_1^{(t)})$$

$$\hat{\beta}_2^{(t+1)} = \hat{\beta}_2^{(t)} + \gamma \nabla Q_2(\hat{\beta}_2^{(t)})$$

$$\nabla Q_k = \begin{bmatrix} \frac{\partial Q_k}{\partial \beta_{0k}} \\ \frac{\partial Q_k}{\partial \beta_{1k}} \\ \frac{\partial Q_k}{\partial \beta_{2k}} \end{bmatrix}$$

5. Multi-Layer deep Neural Nets

Neural Nets can have more than one “hidden” layer.
The outputs from one layer become the inputs to the next.

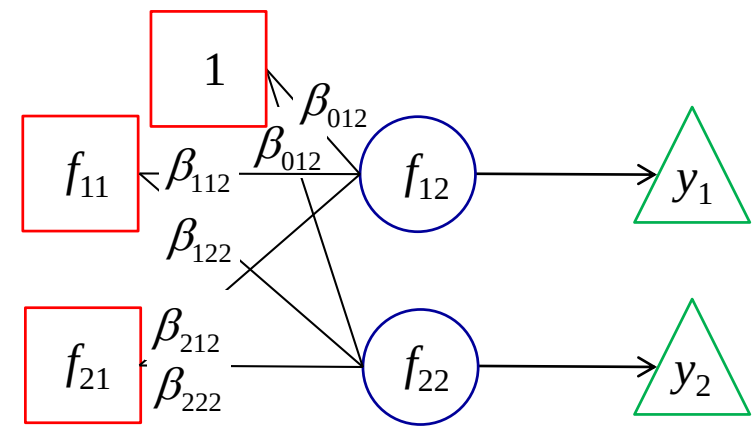


Let’s go through the process as multiple one layer Neural Nets,
from right to left, Backpropagation.

coefficient node layer
 $j = 0,1,2$ $k = 1,2$ $\ell = 1,2$

5. Multi-Layer deep Neural Nets

Neural Nets can have more than one “hidden” layer.
The outputs from one layer become the inputs to the next.

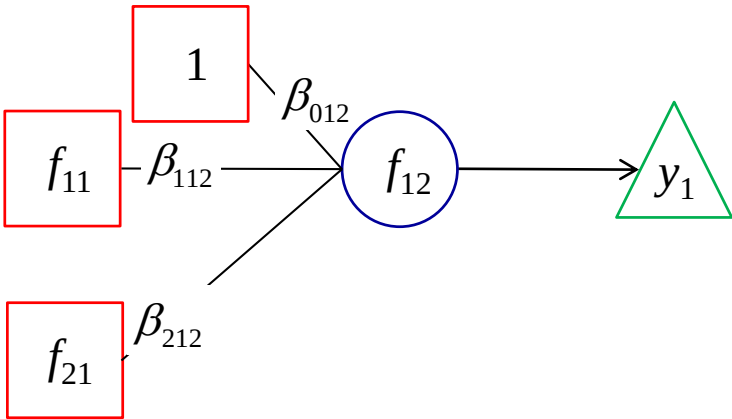


We consider the first output layer as input to the second layer.
Estimate coefficients.

coefficient	node	layer
$j = 0,1,2$	$k = 1,2$	$\ell = 1,2$

5. Multi-Layer deep Neural Nets

Neural Nets can have more than one “hidden” layer.
The outputs from one layer become the inputs to the next.

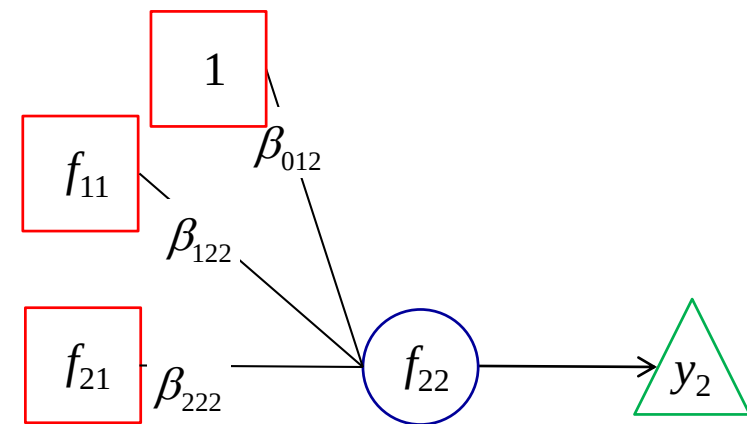


And first focus only on the first node.
Estimate coefficients.

coefficient node layer
 $j = 0,1,2$ $k = 1,2$ $\ell = 1,2$

5. Multi-Layer deep Neural Nets

Neural Nets can have more than one “hidden” layer.
The outputs from one layer become the inputs to the next.

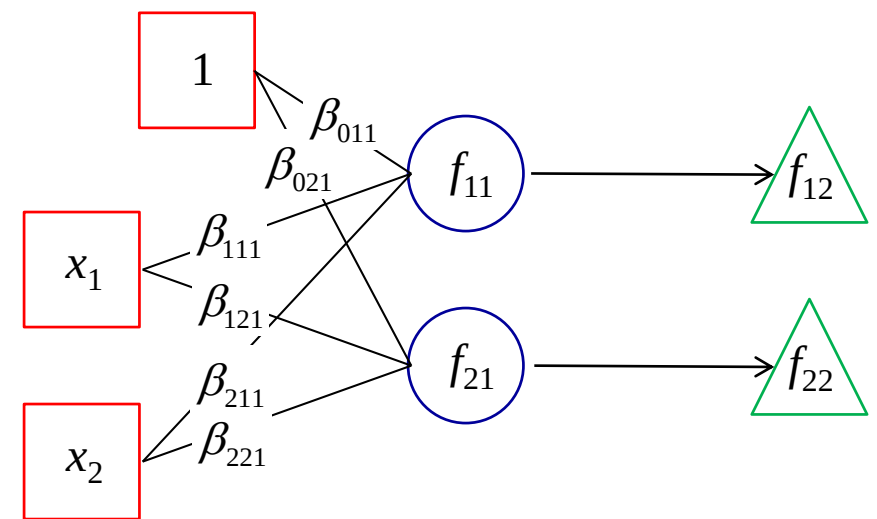


Then focus on the second node.
Estimate coefficients.

coefficient	node	layer
$j = 0,1,2$	$k = 1,2$	$\ell = 1,2$

5. Multi-Layer deep Neural Nets

Neural Nets can have more than one “hidden” layer.
The outputs from one layer become the inputs to the next.

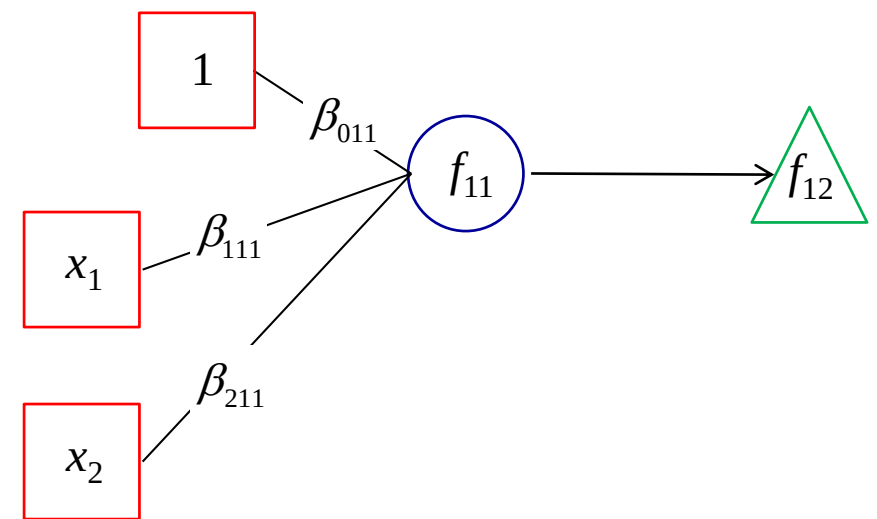


Let’s go through the process as multiple one layer Neural Nets,
from right to left, Backpropagation.

coefficient node layer
 $j = 0,1,2$ $k = 1,2$ $\ell = 1,2$

5. Multi-Layer deep Neural Nets

Neural Nets can have more than one “hidden” layer.
The outputs from one layer become the inputs to the next.

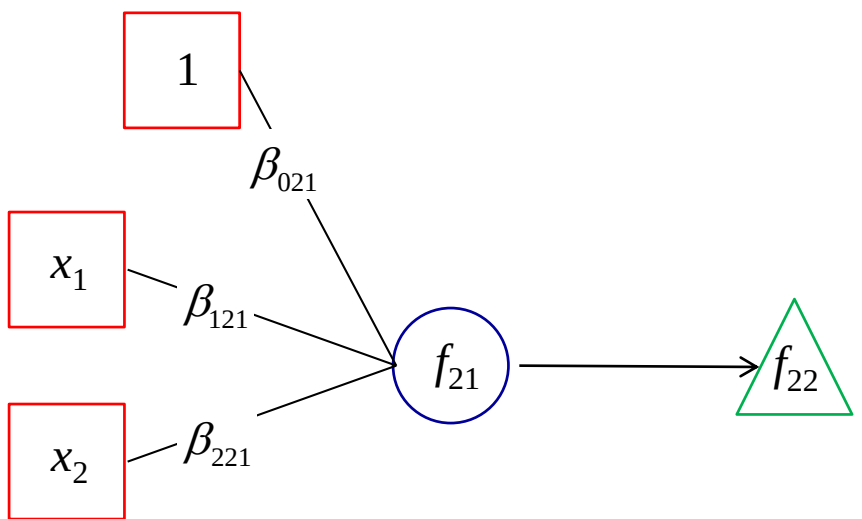


And first focus only on the first node.
Estimate coefficients.

coefficient node layer
 $j = 0,1,2$ $k = 1,2$ $\ell = 1,2$

5. Multi-Layer deep Neural Nets

Neural Nets can have more than one “hidden” layer.
The outputs from one layer become the inputs to the next.

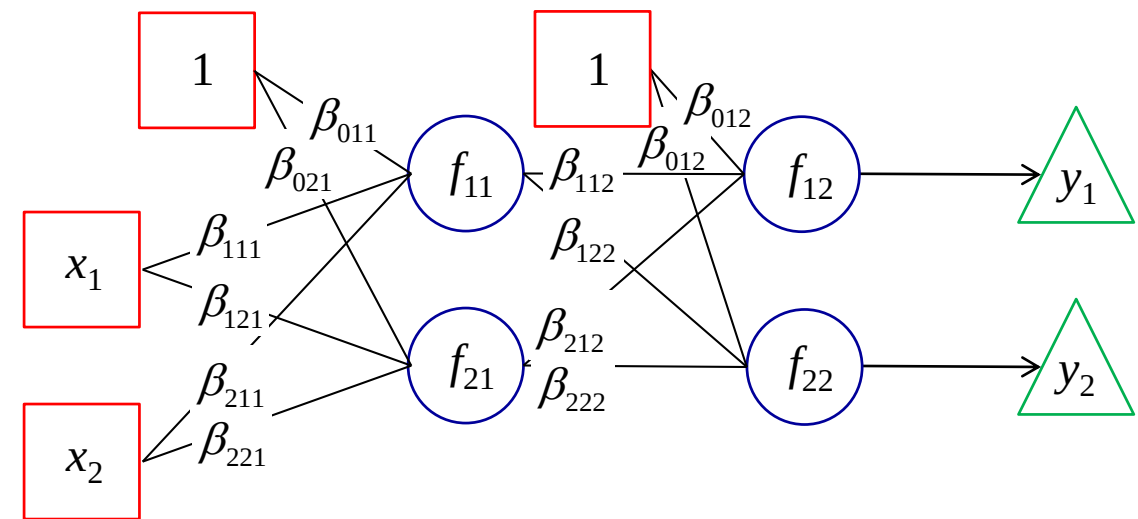


Then focus on the second node.
Estimate coefficients.

coefficient	node	layer
$j = 0,1,2$	$k = 1,2$	$\ell = 1,2$

5. Multi-Layer deep Neural Nets

Neural Nets can have more than one “hidden” layer.
The outputs from one layer become the inputs to the next.

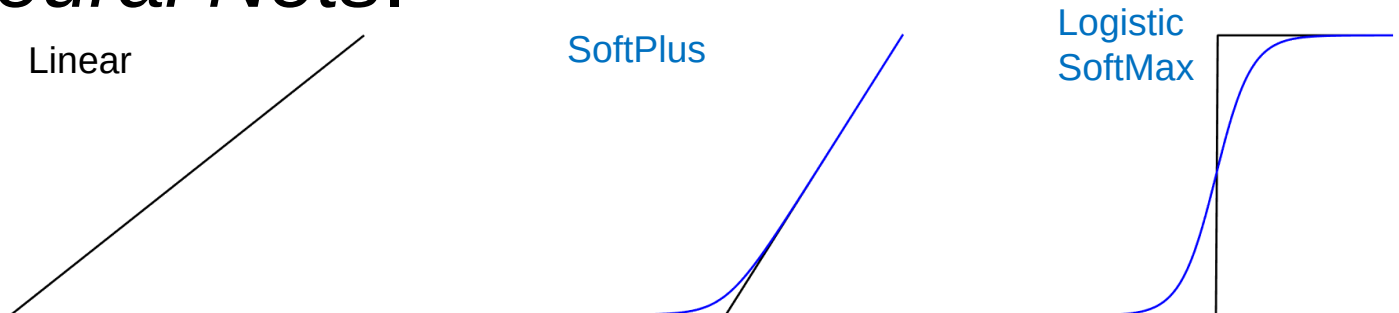


And hence solve the two or multi layer problem.

coefficient node layer
 $j = 0,1,2$ $k = 1,2$ $\ell = 1,2$

5. Discussion

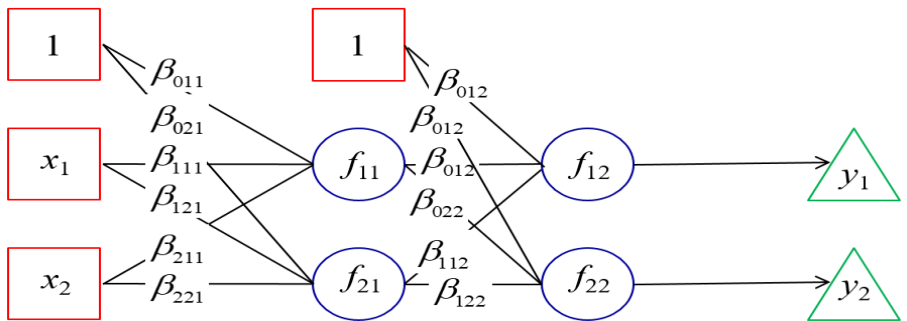
Linear, Logistic, and Non-Linear Regression can be represented as Neural Nets.



Coefficients are estimated via Gradient Descent.

$$\nabla Q = \left(\frac{\partial Q}{\partial \beta_j} \right) \quad \hat{\beta}^{(t+1)} = \hat{\beta}^{(t)} - \gamma \nabla Q(\hat{\beta}^{(t)})$$

Discussed foundational ideas of Neural Nets.
These ideas can be expanded in many directions.



Thank You

Questions?

Daniel.Rowe@Marquette.Edu

MARQUETTE UNIVERSITY
BE THE DIFFERENCE.

MATH 4931/MSSC 5931

— SUMMER 2023

TOPICS IN MATHEMATICS OR STATISTICS: Statistical Machine Vision

MoTuWeTh 11:30am-1:00pm

Topics:

- Discrete image representation.
- Time Series & Image Convolution
- Image enhancement via local pixel weighting (spatial kernel filter and image space convolution).
- Kernel filter design with weight assignments.
- Pixel noise reduction via local averaging (smoothing filters).
- Edge enhancement via local differencing (gradient filters).
- Statistical properties of local averaging or differencing (pixel mean, variance, and correlation).
- Image text recognition, letter or word identification (letter A, word MATH).
- Time averaging (temporal recursive filters) for pixel noise reduction in image sequences.
- Identifying and tracking of objects including orientation through a sequence of images (car moving across a scene in a sequence of images).
- The DFT for accelerated convolutions in frequency space
- Line tracing within an image via discrete derivatives, gradients, and Hessians.
- Image object representations (perimeter, area, elongation, etc.), feature extraction.
- Statistical classification of image objects using features (square, circle, and rectangle).
- Computational implementations and examples will be given with Matlab.
- Additional topics covered if time permitting.

Prerequisites/Notes:

- COSC 1010, MATH 1451, and MATH 4720 or the equiv.
- MSSC 5931 will have additional assignments.

For more information, email the instructor:

- [Dr. Daniel Rowe \(Daniel.Rowe@Marquette.Edu\)](mailto:Daniel.Rowe@Marquette.Edu)

Find text pictures

"If your pictures aren't good enough, you're not close enough."
— Robert Capa

Find lines

Find object

Find industrial concavities

$D > 0$ & $f_{xx}(x_0, y_0) > 0$, then local min at (x_0, y_0) .
 $D > 0$ & $f_{xx}(x_0, y_0) < 0$, then local max at (x_0, y_0) .
 $D < 0$, then f has a saddle point at (x_0, y_0) .
If $D = 0$, anything can happen at (x_0, y_0) .

MARQUETTE UNIVERSITY
BE THE DIFFERENCE.

MATH 4790/MSSC 5790 — FALL 2023

Bayesian Statistics

TuTh 3:30-4:45pm

Tentative topics:

- Conditional probability and Bayes' rule.
- Discrete and continuous distributions of data;
 - binomial, beta, gamma, inverse gamma, normal and Student-t.
 - The bivariate normal and the normal-inverse gamma distributions.
- Maximum likelihood estimation.
- Conditional and marginal distributions.
- Conjugate and non-conjugate prior distributions.
- Maximum a-posteriori and marginal mean estimation.
- Bayesian binomial probability
 - (binomial likelihood, beta prior, beta posterior).
- Bayesian normal mean estimation
 - (normal likelihood, normal-inverse gamma prior, Student-t marginal posterior).
- Bayesian multiple regression
 - (normal likelihood, bivariate normal-inverse gamma prior, bivariate Student-t marginal posterior).
- LASSO (normal likelihood, Laplace-inverse gamma priors).
- Naïve Bayesian Classification
 - (normal class likelihoods, normal-inverse gamma class priors with discrete uniform prior class probabilities).
- Markov chain Monte Carlo numerical integration including
 - importance sampling, Gibbs sampling, and the Hastings algorithm.
- Sequential updating of previous Bayesian models.
- A computational flavor throughout.

Prerequisites/Notes:

- COSC 1010, MATH 1451, and MATH 4720 or the equivalents
- Students enrolled in MSSC 5790 will have additional assignments

For more information, email the instructor:

- [Dr. Daniel Rowe \(Daniel.Rowe@Marquette.Edu\)](mailto:Daniel.Rowe@Marquette.Edu)

$$P(B_i | A) = \frac{P(A | B_i)P(B_i)}{P(A)}$$

$$f(\mu, \sigma^2 | y_1, \dots, y_n) = \frac{f(y_1, \dots, y_n | \mu, \sigma^2) f(\mu, \sigma^2)}{f(y_1, \dots, y_n)}$$