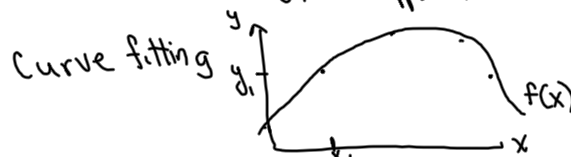


Dr. Cynthia Rudin (CCLS, Columbia)

Regression (the Statistical Learning Version anyway):

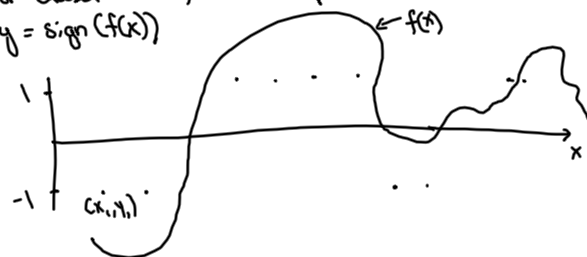
Given $\{(x_i, y_i)\}_{i=1, \dots, m}$ where $x_i \in X$ & $y_i \in \mathbb{R}$, chosen randomly from an unknown probability distribution, find $f \in \mathcal{F}$ $f: X \rightarrow \mathbb{R}$, such that for a new randomly chosen instance $x \in X$, we have $f(x)$ close to its label y .

in some sense which varies on the application



Compare Classification & Regression

For classification, want to produce $f(x)$ such that $y = \text{sign}(f(x))$



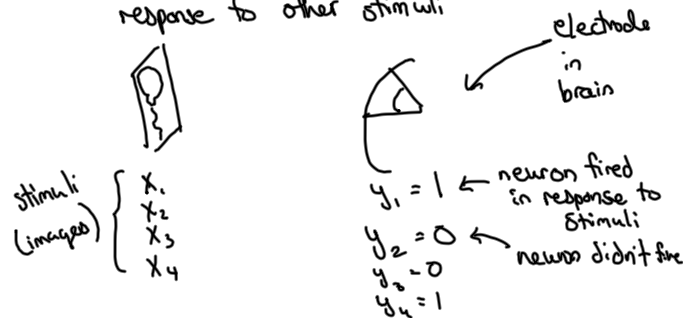
For regression, want to produce $f(x)$ such that $y = f(x)$

Examples:

- 1) Want to estimate the number of customers for a given product, given data about the product and data about past product sales.
- 2) Time Series Analysis, for instance, predicting the price of stocks, or predicting the path of a moving object
- 3) Cancer Genomics (Hans & West 06)
Use gene expression levels (microarray data) from tumor samples to estimate:
 - survival time
 - amount of lymph node invasion
 - response to a particular treatment

Jan 31-2:36 PM

4) Computational Neuroscience: Calculating the firing rate of a neuron in the visual cortex for a new stimuli, based on the neuron's response to other stimuli:



firing rate: prob that neuron will fire in response to stimulus x .

$$= \mathbb{P}(y=1 | x) = \mathbb{E}(y | x) = \int y d\mu(x, y)$$

(# of dimensions = # of pixels of image)

Square Loss and the Mean:

We generally would like to choose $f(x)$ so that it is close to y in some sense, for instance, we might want to choose f to minimize:

$$\mathbb{E}[(y - f(x))^2 | x] \text{ for each } x$$

Say $y \in \{0, 1\}$. Fix x . What is the best $h = f(x)$? Let $p = \mathbb{P}(y=1 | x)$. Turns out the answer is $h = p$, the mean.

We cannot minimize $\mathbb{E}[(y - f(x))^2]$ in practice, because we only have a finite sample of the data $\{(x_i, y_i)\}_{i=1, \dots, m}$.

In practice, we minimize an "empirical" (based on data) quantity:

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

<http://www1.cs.columbia.edu/~allen/S08/NOTES/linkedList1.pdf>

Legendre published & named least squares in 1805 but Gauss claimed he had known the method since 1795. They both used it to determine from astronomical observation the orbits of bodies around the sun.

Solving Least Sq & Pseudo Inverse

Choose f to be linear. $f(\vec{x}) = \vec{w} \cdot \vec{x}$ for $\vec{x} \in \mathbb{R}^n$

$$\text{Least Sq Problem: } \min_{\vec{w} \in \mathbb{R}^n} \frac{1}{m} \sum_{i=1}^m (\vec{w} \cdot \vec{x}_i - y_i)^2$$

Least $\mathcal{J}_g$: $\min_{\vec{w} \in \mathbb{R}^n} \underbrace{\frac{1}{m} \sum_{i=1}^m (\vec{w} \cdot \vec{x}_i - y_i)^2}_{\Phi(\vec{w})}$

To solve, set $\nabla \Phi(\vec{w}) = 0$

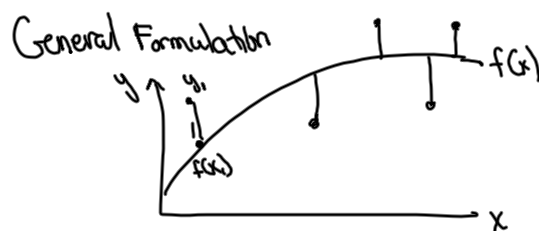
$$\nabla \Phi(\vec{w}) = \begin{pmatrix} \partial \Phi / \partial w_1 \\ \vdots \\ \partial \Phi / \partial w_n \end{pmatrix} = \frac{1}{m} 2 X^T (X \cdot \vec{w} - \vec{y})$$

where $X = \begin{pmatrix} \text{---} & x_1^T & \text{---} \\ & \vdots & \\ \text{---} & x_n^T & \text{---} \end{pmatrix}$

$$\nabla \Phi(\vec{w}) = 0 \Rightarrow \vec{w} = \underbrace{(X^T X)^{-1}}_{\text{pseudoinverse}} X^T \vec{y}$$

If columns of X are indep., pseudoinverse exists.
If X has lin. indep. columns & rows, X^{-1} is the pseudoinverse

Pseudoinverse: Fredholm 1903
Moore 1920, Penrose 1955



For regression (& other learning problems) we often try to minimize a function of the form

$$\mathcal{L}(f) = \underbrace{\sum_{i=1}^m \ell(y_i, f(x_i))}_{\text{"empirical error" term measures closeness of data to model}} + \underbrace{c R(f)}_{\substack{\text{regularization} \\ \text{parameter } c \in \mathbb{R}}} \underbrace{\text{regularization term measures smoothness of } f \text{ to prevent overfitting}}$$

If c is too small, model may overfit

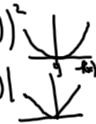
If c is too large, " " underfit



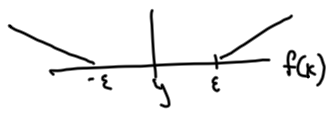
If c is just right, model generalizes (predicts) well



Some loss functions for regression:

least squares (l_2): $l(y, f(x)) := (y - f(x))^2$ 

absolute distance (l_1): $l(y, f(x)) := |y - f(x)|$ 

SVM ϵ -insensitive $l(y, f(x)) := \begin{cases} 0 & \text{if } |y - f(x)| \leq \epsilon \\ |y - f(x)| & \text{otherwise} \end{cases}$ 

Some regularization terms, assuming $f_w(\vec{x}) = \vec{w} \cdot \vec{x}$:

$$R(f_w) = \|\vec{w}\|_2^2 = \sum_j w_j^2$$

$$R(f_w) = \|\vec{w}\|_1 = \sum_j |w_j|$$

$$R(f_w) = \|\vec{w}\|_{\text{RKHS}}^2 \quad (\text{you'll hear plenty about this later!})$$

Tikhonov Regularization Problem or Ridge Regression:

$\min_{\vec{w}} \|\vec{X}\vec{w} - \vec{y}\|_2^2 + \|\vec{\Gamma}\vec{w}\|_2^2$ can be solved:
 $\vec{\Gamma}$ pos. semidef, $\vec{z} \cdot \vec{\Gamma} \cdot \vec{z} \geq 0 \quad \forall \vec{z} \in \mathbb{R}^n$

$$\vec{w} = (\vec{X}^T \vec{X} + \vec{\Gamma}^T \vec{\Gamma})^{-1} \vec{X}^T \vec{y}$$

True l_0 sparsity

There has been a recent flurry of work (inspired by a result of Donoho) showing that in some cases (when columns of $\vec{X}$ are almost orthogonal) the minimizer of:

$$(*) \quad \|\vec{X}\vec{w} - \vec{y}\|_2^2 + C \|\vec{w}\|_1$$

is the same as the minimizer of the real sparsity problem:

$$(*2) \quad \|\vec{X}\vec{w} - \vec{y}\|_2^2 + C \|\vec{w}\|_0$$

where $\|\vec{w}\|_0$ = the number of nonzero elements of $\vec{w}$.

However $(*)$ is convex & we can solve it whereas $(*2)$ is not!

Linear vs. Nonlinear Modelling

There are many options for the form of the function f .

$$f(\vec{x}) = \vec{w} \cdot \vec{x} \quad \text{linear}$$

$$f(\vec{x}) = \vec{w} \cdot \phi(\vec{x}) \quad \text{"non-linear"}$$

ϕ can be a vector in an infinite-dimensional Hilbert space.

This means $\vec{w}$ is infinite dimensional too!

This intro to regression is only the beginning!
Many other techniques:

- splines
- regression trees
- ~~off~~ networks & other nonlinear methods

In 1900 3 dims was a lot.

In 1980's 10 dims was a lot.

Now, we have $\sim 10,000$ dim