

Examples $\{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\}$

$x_i \in \mathbb{R}^d$, where $d = \#$ of features

$y_i \in \mathbb{R}$ is a label.

training data set

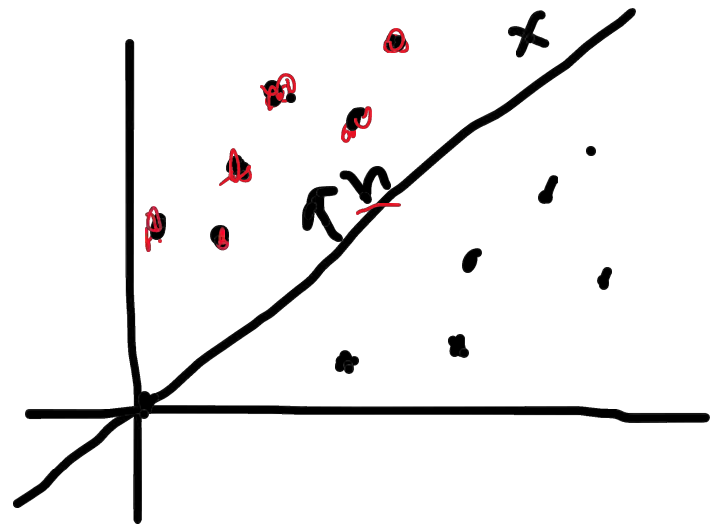
Goal: learn a model h , sometimes called a hypothesis, from some family $\mathcal{H}$ of possible models.

Given a new unlabeled point $x \in \mathbb{R}^d$,
use your model h to predict label y .
This x is called test data.

How well h predicts y is called
generalization error.

Linear Classification:

labels are 0 or 1
red black.



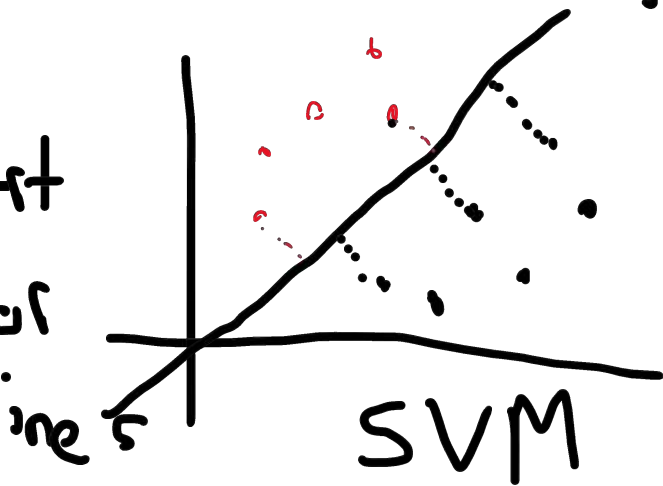
$\langle h, x \rangle$
h is
normal
vector
to
line

Training data: $(x_1, 0), (x_2, 1), (x_3, 1), \dots$

Want h to separate red and black points.

Want: $\langle x_i, \underline{h} \rangle \geq 0$ for all red points x_i
 $\langle x_i, \underline{h} \rangle \leq 0$ for all black points x_i

Linear Classification with a margin γ



categorical - output is binary
which line should I choose?

choose maximum margin line.
closest points called support vectors

Margin: want every point to be distance $\geq \gamma$
from the separating line. How to find h ?

want $\langle x_i, h \rangle \geq \gamma$ for all red points
 $\langle x_i, h \rangle \leq -\gamma$ for all black points

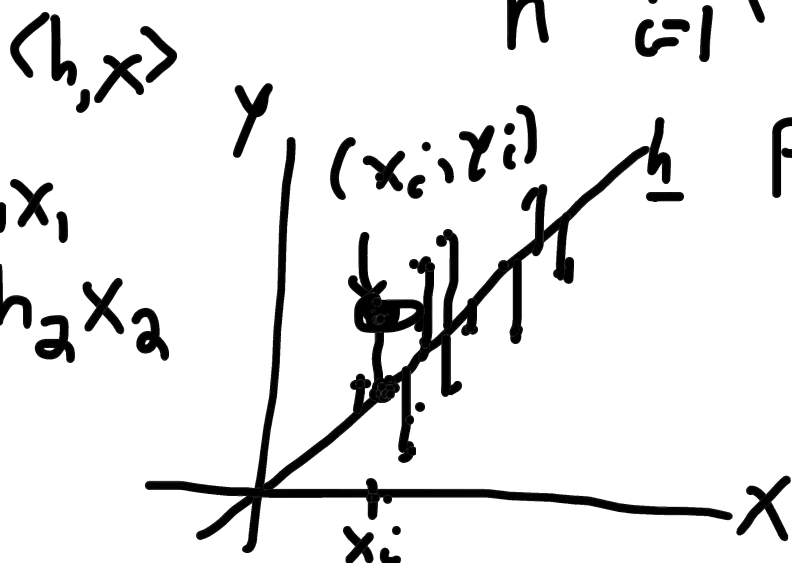
binary search
for γ .

Numerical output - example is linear regression. ^{least squares} test y data.

Training data $(x_1, y_1), \dots, (x_n, y_n)$ but now $y_i \in \mathbb{R}$.

Model h is chosen to minimize loss function $L(h)$:

$$L(h) = \frac{1}{n} \sum_{i=1}^n (\underbrace{\langle x_i, h \rangle}_{\text{prediction}} - \underbrace{y_i}_{\text{label}})^2$$



find a line to "fit" your data
 $x_i \in \mathbb{R}$ in example
(in general $x_i \in \mathbb{R}^d$)

How to solve regression with gradient descent

$$L(h) = \frac{1}{n} \sum_{i=1}^n (\langle x_i, h \rangle - y_i)^2$$

minimize
a convex
function

Initialize h^0 (do it randomly)

For $t = 1, 2, \dots, T$

$$h^t \leftarrow h^{t-1} - \eta_t \nabla L|_{h^{t-1}}$$

← what is the
gradient?

End

Output $\frac{1}{T} \sum_{t=1}^T h^t$ (in practice just output h^T),

Need gradient of $L(h) = \frac{1}{n} \sum_{i=1}^n (\langle h, x_i \rangle - y_i)^2$

Gradient is linear!!

$$\nabla L = \frac{1}{n} \sum_{i=1}^n \nabla (\underbrace{\langle h, x_i \rangle - y_i}_{f_i})^2$$

$$= \frac{1}{n} \sum_i \nabla f_i$$

$x_i \in \mathbb{R}^d$

$h \in \mathbb{R}^d$

$\nabla(f+g) = \nabla f + \nabla g$

$\nabla \circ f = \nabla f$

$$\nabla L = \left(\frac{\partial L}{\partial h_1}, \frac{\partial L}{\partial h_2}, \dots, \frac{\partial L}{\partial h_d} \right)$$

$$\nabla f_i = (2(\langle h, x_i \rangle - y_i) x_{i1}, 2(\langle h, x_i \rangle - y_i) x_{i2}, \dots)$$

$\nabla f_i|_{h^{t-1}}$

$\nabla L|_{h^{t-1}} = \frac{1}{n} \sum_{i=1}^n \nabla f_i|_{h^{t-1}}$

chain rule!

Chain Rule: $f(x) = u(\overbrace{v(x)})$

slope $\downarrow$ evolution
 $y = \underbrace{m \cdot x}_{\text{evolution}}$

Equation of line

line goes through

x is 1-dimensional

$$v(h_1) =$$

$$\frac{\partial f}{\partial x} = \frac{\partial u}{\partial v} \cdot \frac{\partial v}{\partial x}$$

$$(\frac{\partial u}{\partial v}) \cdot x_{ij}$$

$$\langle h, x_i \rangle - y_i$$

$$= \sum_{j=1}^d h_j \cdot x_{ij} - y_i$$

$$h_j \cdot \underline{x_{ij}} \quad u(v(h_1)) = v(h_1)$$

$$2(\langle h, x_i \rangle - y_i) x_{if}$$

$$f_i = (\langle h, x_i \rangle - y_i)^2$$

$$\frac{\partial f_i}{\partial h_1}, \frac{\partial f_i}{\partial h_2}, \dots$$

Overfitting: you don't want to choose a
"very weird" looking x to fit your data.

Occam's Razor: "simplest" solutions generalize better. adj $\lambda > 0$

In regression, don't want h to have large norm.

Ridge regression: $L(h) = \frac{1}{n} \sum_{i=1}^n (\langle h, x_i \rangle - y_i)^2 + \lambda \|h\|^2$ $\lambda > 0$

Solve with gradient descent!

What if n is really large? Computing gradient is slow.

Stochastic Gradient Descent (SGD)

$$\nabla L = \frac{1}{n} \sum_{i=1}^n \nabla f_i \quad \text{slow because } n \text{ is large}$$

Randomly sample $I \in \{1, 2, \dots, n\}$

$$\tilde{\nabla} L = \nabla f_I \quad \text{much faster - just one } \nabla f_I \text{ calculation!}$$

$$E[\tilde{\nabla} L] = \frac{1}{n} \sum_{i=1}^n \nabla f_i = \nabla L \quad \text{so unbiased estimator.}$$

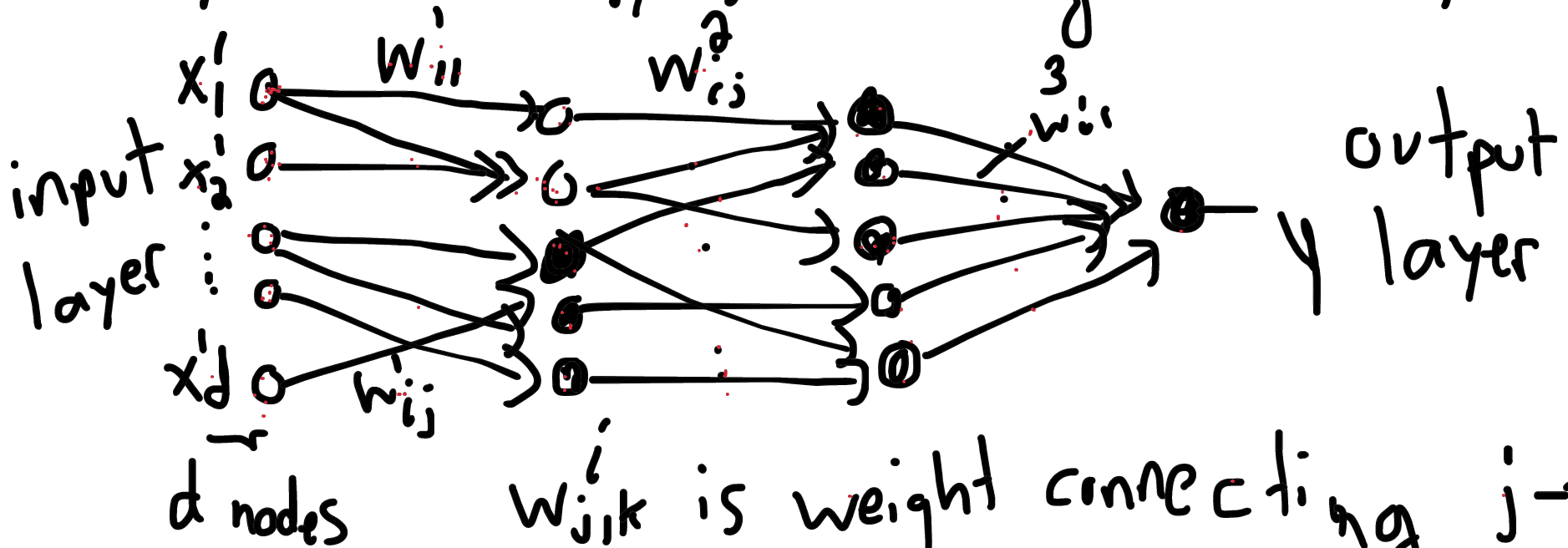
Reduce variance of SGD by using minibatches.
Instead of sampling 1 example, we sample b examples. $b = \text{batch size}$. $I_1, I_2, \dots, I_b$ are sampled examples.

$$\tilde{\nabla} L = \frac{1}{b} \sum_{j=1}^b \nabla f_{I_j}.$$

If $b=n$, get GD. If $b=1$, get SGD.

b provides a time versus variance tradeoff.

Feedforward Neural Networks: $x^i \in \mathbb{R}^d$
 $(x^1, y^1), \dots, (x^n, y^n)$ is training data. $y^i \in \mathbb{R}$.

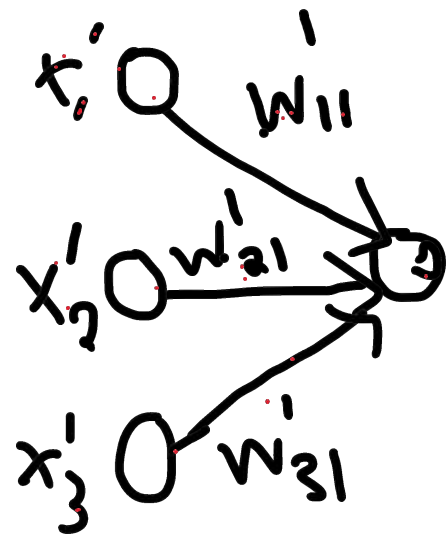


$w_{j,k}^i$ is weight connecting j -th node in i -th layer to k -th node in $i+1$ layer

Zoom in

Zoom In:

How to learn weights?



$$f(\underline{x_1' \cdot w_{11} + x_2' w_{21} + x_3' w_{31}})$$

f is activation
(non-linear) function

ReLU (rectified linear unit)



Sigmoid

