

RMS Prop (2012)

attempted to fix "prob" w/ Adagrad of inevitable learning rate decay

standard to take $\beta = 0.9$

Idea: $D_t = \beta D_{t-1} + (1-\beta) \text{diag} \left(\nabla f_t(\omega_t) \nabla f_t(\omega_t)^T \right)$

vs

$D_t = D_{t-1} + \text{diag} \left(\nabla f_t(\omega_t) \nabla f_t(\omega_t)^T \right)$ for Adagrad

Let $g_t = \nabla f_t(\omega_t)$ be our gradient or stochastic gradient estimate at time t , then

$$\text{diag} (g_t g_t^T) = \begin{bmatrix} (g_t)_1^2 & & \\ & \ddots & \\ & & (g_t)_d^2 \end{bmatrix}$$

e.g.

$$\begin{aligned} D_3 &= .9 D_2 + .1 \text{diag}(g_3 g_3^T) \\ &= .9 (.9 D_1 + .1 \text{diag}(g_2 g_2^T)) + .1 \text{diag}(g_3 g_3^T) \\ &= \underline{(.9)^2} D_1 + \underline{(.9)(.1)} \text{diag}(g_2 g_2^T) \\ &\quad + \underline{.1} \text{diag}(g_3 g_3^T) \end{aligned}$$

This update prevents inevitable learning rate decay

$$\omega_{t+1} = \omega_t - \alpha_t (D_t + \epsilon I)^{-1/2} \nabla f_t(\omega_t)$$

ADAM ^{Ringma} (2014) very popular in deep learnings

- tries to fix the diminishing step size problem of Adagrad
- also tries to estimate a better search direction using momentum (exponential average)

$$\mu_t = \beta_1 \mu_{t-1} + (1 - \beta_1) \nabla f_t(w_t)$$

$$D_t = \beta_2 D_{t-1} + (1 - \beta_2) \text{diag}(\nabla f_t(w_t) \nabla f_t(w_t)^T)$$

$$w_{t+1} = w_t - \alpha_t (D_t + \epsilon I)^{-1/2} \mu_t$$

General approach to nonlinear ML: Kernel Methods

Instead of models of the form $f(x) = \sum_{i=1}^d \omega_i x_i + \omega_0$

(linear models) we want to use/fit

nonlinear models $f(x) = ?$

As our working example, consider quadratic model of the form

$$f(x) = \omega_0 + \sum_{i=1}^d \omega_i x_i + \sum_{i,j=1}^d \omega_{i,j} \underbrace{x_i x_j}_{\text{nonlinear}}$$

Note the parameter "vector" is $\left[\overset{\mathbb{R}}{\omega_0}, \overset{\mathbb{R}^d}{\omega}, \overset{\mathbb{R}^{d^2}}{W} \right] \leftarrow \underbrace{1}_{1} + \underbrace{d}_{d} + \underbrace{d + \binom{d}{2}}_{d^2} \text{ (quadratic) interaction of features}$

so we have $\mathcal{O}(d^2)$ parameters

note $f(x) = \omega_0 + \langle \omega, x \rangle + x^T W x$

Consider the case of using $f(x)$, a quadratic model, for regression. We use LS loss and ℓ_2 regularization.

Flatten our model parameters into a vector $v \in \mathbb{R}^{1+2d+\binom{d}{2}}$ and write

$$f(x) = v^T \begin{bmatrix} 1 \\ x_1 \\ \vdots \\ x_d \\ x_1^2 \\ x_1 x_2 \\ \vdots \\ x_d^2 \end{bmatrix} = v_0 + \sum_{i=1}^d v_i x_i + \sum_{i,j=1}^d v_{d+(i-1)d+j} x_i x_j$$

so $v_0 = \omega_0$

$v_{1:d} = \omega$

$v_{(d+1):(1+2d+\binom{d}{2})} \stackrel{\text{reshape}}{=} W$

$= v^T \phi(x)$

where $\phi: \mathbb{R}^d \rightarrow \mathbb{R}^{1+d+\binom{d}{2}}$

$$\phi\left(\begin{bmatrix} x_1 \\ \vdots \\ x_d \end{bmatrix}\right) = \begin{bmatrix} 1 \\ x_1 \\ \vdots \\ x_d \\ x_1^2 \\ x_1 x_2 \\ \vdots \\ x_d^2 \end{bmatrix}$$

is called our feature map

and we want to solve the ERM problem

$$\begin{aligned} v_* &= \arg\min_v \frac{1}{n} \sum_{i=1}^n (y_i - f(x_i))^2 + \lambda \|v_{1:d}\|_2^2 \\ &= \arg\min_v \frac{1}{n} \sum_{i=1}^n (y_i - v^T \phi(x_i))^2 + \lambda \|v_{1:d}\|_2^2 \end{aligned}$$

don't penalize the bias

v_0
↓

Define feature matrix

$$\bar{\Phi} = \begin{bmatrix} \phi(x_1)^T \\ \vdots \\ \phi(x_n)^T \end{bmatrix} \in \mathbb{R}^{n \times D}$$

$$\rightarrow D = 1 + d + \binom{d}{2}$$

then we have

$$v_* = \operatorname{argmin} \frac{1}{n} \|y - \bar{\Phi} v\|_2^2 + \lambda \|v_{1:}\|_2^2$$

$$v_* = \underbrace{(\bar{\Phi}^T \bar{\Phi} + n\lambda \mathbf{I})}^{D \times D = \mathcal{JZ}(d^Z \times d^Z)}^{-1} \bar{\Phi}^T y \leftarrow \left(\begin{array}{l} \text{assuming we} \\ \text{regularized w/} \\ \lambda \|v\|_2^2 \end{array} \right)$$

Now we can predict

$$f(x_{\text{new}}) = v_*^T \phi(x_{\text{new}})$$

drawback: using Φ and solving the linear system
requires inverting a $D \times D$ matrix which is infeasible
if $d > 10K$ $\hat{=}$ " $d^2 \times d^2$ "

remedy: I only care about prediction; we can do
this much faster.

$$f(x_{\text{new}}) = v_*^T \phi(x_{\text{new}}) = \phi(x_{\text{new}})^T (\Phi^T \Phi + n\lambda I)^{-1} \Phi^T y$$

observe (from the SVD) that for any matrix M

$$(M^T M + n\lambda I)^{-1} M^T = M^T (M M^T + n\lambda I)^{-1}$$

$$\Rightarrow f(x_{\text{new}}) = \phi(x_{\text{new}})^T \underbrace{\Phi^T (\Phi \Phi^T + n\lambda I)^{-1}}_{n \times n} y$$

Implies prediction is cheaper than the naive approach
when $n \ll D$

Let's introduce the kernel function $\overset{\text{lowercase}}{\kappa}(x, y)$
where $x, y \in \mathbb{R}^d$

$$\kappa(x, y) = \langle \phi(x), \phi(y) \rangle$$

associated w/ the feature map ϕ

Further introduce the kernel matrix on our training data

$\overset{\text{uppercase}}{\text{letter}} \kappa \left[\begin{matrix} \kappa \end{matrix} \right]_{ij} = \kappa(x_i, x_j) \quad \text{where} \quad \kappa \in \mathbb{R}^{n \times n}$

Now note that with this notation,

$$\begin{aligned}\phi(x_{\text{new}})^T \bar{\Phi}^T &= \phi(x_{\text{new}})^T \begin{bmatrix} \phi(x_1) \\ \vdots \\ \vdots \\ \phi(x_n) \end{bmatrix} \\ &= \begin{bmatrix} K(x_{\text{new}}, x_1) & \dots & K(x_{\text{new}}, x_n) \end{bmatrix}\end{aligned}$$

and

$$(\bar{\Phi} \bar{\Phi}^T + n \lambda \mathbf{I})^{-1} y$$

$$\uparrow (\bar{\Phi} \bar{\Phi}^T)_{ij} = \langle \phi(x_i), \phi(x_j) \rangle = K(x_i, x_j)$$

$$\Leftrightarrow \bar{\Phi} \bar{\Phi}^T = K$$

so

$$f(x_{\text{new}}) = \begin{bmatrix} K(x_{\text{new}}, x_1) & \dots & K(x_{\text{new}}, x_n) \end{bmatrix} (K + n \lambda \mathbf{I})^{-1} y$$

Letting $\alpha = (K + n\lambda I)^{-1}y$ be precomputed, we can predict by forming

$$f(x_{\text{new}}) = \sum_{i=1}^n K(x_{\text{new}}, x_i) \alpha_i$$

we call f a kernel estimator

Comparison
for linear models $f(x_{\text{new}}) = \langle \alpha, \underbrace{x_{\text{new}}}_{\mathbb{R}^d} \rangle$

for nonlinear models, our feature map ϕ induces a kernel K and our predictions take the form

$$f(x_{\text{new}}) = \left\langle \underbrace{\begin{bmatrix} K(x_{\text{new}}, x_1) \\ \vdots \\ K(x_{\text{new}}, x_n) \end{bmatrix}}_{\mathbb{R}^n \rightarrow \mathbb{R}^n}, \alpha \right\rangle$$

We presented this result for a specific
loss function and model class
(least squares) (quadratic)

In general we can use the kernel trick idea to kernelize
machine learning problems for a large class of
hypothesis spaces w/ specific regularizers

Kernel Methods

(RKHS)

$\mathcal{H}$ needs to be a reproducing kernel Hilbert space
and the regularizer $R(f) = \|f\|_{\mathcal{H}}^2$ where $\|\cdot\|_{\mathcal{H}}$
is the Hilbert space norm.

Then the ERM

$$f^* = \underset{f \in \mathcal{H}}{\operatorname{argmin}} \quad \frac{1}{n} \sum_{i=1}^n \ell(f(x_i), y_i) + \frac{\lambda}{2} \|f\|_{\mathcal{H}}^2$$

has a solution of the form

$$f^*(x) = \sum_{i=1}^n K(x, x_i) \alpha_i$$

where K is the kernel associated with $\mathcal{H}$ ($K \in \mathcal{H}$
determine
each other)