

Extreme Gradient Boosting

Can be used for classification & regression. Takes in multiple input features (x) to predict a target (y).
Let's discuss the case of XGB Regressor.

e.g.

row# (i)	Day of Week	Price	Promotion	Demand
Day	x1	x2	x3	y
1	1	10	0	100
2	2	9	1	120
3	3	11	0	90
4	4	10	1	130
5	5	12	0	85

①

Initialize Prediction

Usually called base prediction

Mean of targets

$$\hat{y}_i^{(0)} = \bar{y} = \frac{100 + 120 + 90 + 130 + 85}{5} = 105$$

Base
Prediction for i^{th} row

②

Squared Error Loss

$$L = \frac{1}{2} (y_i - \hat{y}_i)^2$$

Loss function
Actual
Prediction

Gradient

$$g_i = - \frac{\partial L}{\partial \hat{y}_i} = y_i - \hat{y}_i$$

$$g_i = \hat{y}_i - y_i$$

Prediction
Actual

$$h_i = \frac{\partial^2 L}{\partial \hat{y}_i^2} = \frac{\partial}{\partial \hat{y}_i} \left(\frac{\partial L}{\partial \hat{y}_i} \right) = \frac{\partial}{\partial \hat{y}_i} (y_i - \hat{y}_i) = -1$$

e.g.

i	Actual y_i	Actual gradient $(\hat{y}_i - y_i)$ $(105 - y_i)$
1	100	5
2	120	-15
3	90	15
4	130	-25
5	85	20

The goal is to improve prediction $\hat{y}_i$ at every step by adding a corrective factor

$$\hat{y}_i^{(t)} = \hat{y}_i^{(t-1)} + f_i(x)$$

Prediction at step t
Prediction at last step
Corrective term

$$\hat{y}_i^t = \hat{y}_i^{t-1} + \eta w_i$$

That is why we use gradient - To see where we going (steepest descent like gradient descent)

③ Tree

<u>Left</u>	<u>Right</u>
Leaf 1	Leaf 2
Day 1-2	Day 3-5
$G_L = \sum_L g_i = -10$	$G_R = \sum_R g_i = 10$
$H_L = \sum_L h_i = 2$	$H_R = \sum_R h_i = 3$

We measure how good this split is, for which we compute 'Gain'

$$\text{Gain} = \frac{1}{2} \left[\frac{G_L^2}{H_L + \lambda} + \frac{G_R^2}{H_R + \lambda} - \frac{(G_L + G_R)^2}{H_L + H_R + \lambda} \right] - \gamma$$

$\rightarrow$ Regularization
 $\rightarrow$ Penalty term for split
 If set to low - splits are favored & can lead to complex trees
 " high - simple trees

Split with most Gain is preferred

$$\omega_j = \frac{\sum_L g_i}{\sum_L h_i + \lambda}$$

$\rightarrow$ If > 0 underprediction
 < 0 overprediction
 $\rightarrow$ Regularization
 $\rightarrow$ Total curvature ($\because$ 2nd derivative)
 Confidence of g
 Set of samples in leaf

Going to the right leaf

Right
Leaf 2

Day 3-5

$$\omega_j = \frac{10}{4} = 2.5$$

$G_R = \sum_R g_i = 10$
 $H_R = \sum_R h_i = 3$

$\therefore$ all samples in this leaf would have prediction decreased by 2.25 ($\lambda = 0$)

Left
Leaf 1

Day 1-2

$$\omega_j = \frac{-10}{3} = -3.33$$

$G_L = \sum_L g_i = -10$
 $H_L = \sum_L h_i = 2$

④ Update Predictions

$$\hat{y}_i^{(1)} = \hat{y}_i^{(0)} + \eta \omega_j$$

$\rightarrow$ Learning Rate
 $\rightarrow$ Weight
 Prediction in next iteration (First)
 Prediction previous iteration (Base)

Right

$$\hat{y}_i^{(1)} = \hat{y}_i^{(0)} + \eta \omega_j$$

$$y_i^{(1)} = 105 + (0.1)(2.5) = 105.25$$

Left

$$y_i^{(1)} = 105 + (0.1)(-3.33)$$

$$= 104.667$$

Repeat

$$\hat{y}_i = \hat{y}_i^{(0)} + \sum_{t=1}^T \eta f_t(x_i)$$

Actual eq. with Regularization

Objective Function

$$\mathcal{L}^{(t)} \approx \sum_{i=1}^n \left[g_i(f_t(x_i)) + \frac{1}{2} h_i(f_t(x_i)^2) \right] + \Omega(f_t)$$

Annotations:

- $g_i(f_t(x_i))$: Gradient of loss w.r.t $t-1$
- $h_i(f_t(x_i)^2)$: Hessian
- $\Omega(f_t)$: Regularization
- Output of new ω_j

Taylor Approx. $\therefore$

$$g_i = \frac{\partial \mathcal{L}}{\partial y_i^{t-1}}, h_i = \frac{\partial^2 \mathcal{L}}{(\partial y_i^{t-1})^2}$$

So, it's in the form of Taylor series
The advantage of using this Taylor series approx.
is you don't need to define exact loss function