

Chapter 7

How to select K in kNN (1Q)

K= 1: By validation data; whichever k gives the lowest validation error.

Binary Classification K With Even K's

DO NOT USE even numbers since it could lead to a tie. XLMiner will pick the lowest probability and can choose an even number but that doesn't mean it should be chosen

K > 1: classify by the majority decision rule based on the nearest k records

Low K values:

Capture local structure but may also capture noise. You can't rely on one Neighbour

High K values:

Provide more smoothing but may lose local detail. K can be as large as the training sample

Choose the K that gives you the lowest valid ER

Euclidean Distance (1Q)

sometimes predictors need to be standardized to equalize scales before computing distances. Standardized = normalized (-3, 3)

of possible partitions in Recursive Partition (2Q)

Continuous: $(n-1)*p$

Categorical: 2ID - 1P, 3ID - 3P, 4ID - 7P, 5ID - 15P

Cut Off Value in Classification

- Cutoff = 0.5 by default because the proportion of observation neighbors 1's in the k nearest neighbors. Majority decision rule is related to the cut off value for classifying records

You can adjust the cut off value to improve accuracy

Y = 1 (if $p > \text{cutoff}$)

Y = 0 (if $p < \text{cutoff}$)

Cut Off Example Question

Example: Suppose cutoff = 0.9, k=7, we observed 5 C1 and 2C0. Y = 1 or 0?

- Probability ($Y=1$) = $5/7 = 0.71 \rightarrow 0.71 < 0.9 \rightarrow Y = 0$

Regression Tree

- Used with **continuous** outcome variables. Many splits attempted, choose the one that minimizes impurity

- Prediction is computed as the **average** of numerical target variables in the rectangle

- **Impurity measured by the sum of squared deviation from leaf mean**

- Performance measured by RMSE

Regression Tree is used for prediction. Compared to classification tree, we only have to ...

Replace impurity measure by the sum of squared deviation everything else will be the same.

Split by irrelevant variables = Bad impurity score

Only split with relevant variables

General Info

- Makes no assumptions about the data
- Gets classified as whatever the predominant class is among nearby records
- the way to find the k nearest neighbors in Knn is through the Euclidean distance

Rescaling: Only for kNN do you need to rescale because the amount of contribution from each variable. No need for logistic regression since it does not change the P value or RMSE

No need for CART since it doesn't change the order of values in a variable

XLMiner can only handle up to K= 10

Chapter 9

Properties of CART (Classification And Regression Tree)(3Q)

- Model Free

- Automatic variable selection

- Needs large sample size (bc its model free)

- Only gives horizontal or vertical splits

- both methods of CART are BOTH model free

Best pruned tree: You naturally get overfitting when the natural end of process is 100% purity in each leaf which ends up fitting noise in the data. Slightly overfitted so people partition a bit less to accommodate based on the minimal error tree

Minimum error tree: The tree with lowest validation error.

Full tree: largest tree training error equals zero; overfitted



By angelica9373

cheatography.com/angelica9373/

Published 19th October, 2024.

Last updated 22nd October, 2024.

Page 1 of 3.

Sponsored by **CrosswordCheats.com**

Learn to solve cryptic crosswords!

<http://crosswordcheats.com>

Chapter 9 (cont)

Note: The full tree can be the same as the minimum error tree BUT usually best pruned tree should be smaller than the other trees

Impurity Score

- Metric to determine the homogeneity of the resulting subgroups of observations
- For both, the lower the better
- One has no advantage using one over the other.

Gini Index (0, 0.50 binary)

Entropy Measure: (0, $\log_2^2$ if binary) OR (0, $\log_2(m)$ -> m is the total # of classes of Y)

Overall Impurity Measure: Weighted average of impurity from individuals' rectangles weights being the proportion of cases in each rectangle.
Choose the split that reduces impurity the most (split points becomes nodes on the tree)

Check notes for that distance = to weighted average ratio

Dimensional Predictors Q's

Continuous Partitions

- (n-1) x P -> p dimensional predictors (more than 2 dimensional predictors)

Categorical Partitions

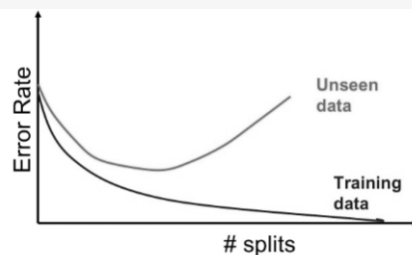
- abcd split. (3Levels, 3P), (4Levels, 7P)

XLMiner only supports binary categorical variables

When to Stop Partitioning

- Error rate as a function of the number of splits for training vs validation data -> Indicates overfitting
- We can continue partitioning the tree so a **FULL tree** will be obtained in the end. A full tree is usually overfitted so we have to impose and **EARLY STOP** ...
- Stop when training error rate is approaching 0 as you partition further but you must have an early stop before letting it touch 0. **Early - Stop** (Minimum Error Tree or Best Prune tree): OR Stop based off **Chi-square tests**:
- if the improvement of the additional split is **statistically significant** -> continue. If not, STOP.
- Largest to Smallest:** Full Tree > Min error tree > Best prune tree (Std usually smaller than min error). *Keep in mind: Full tree CAN BE THE SAME as your Min error tree*

Error Rate as you continue Splitting



Training error decreases as you split more
Validation error decreases and then increases with the tree size

Chapter 10

Assumptions For Logistic Regression (1Q)

- Generalized linearity

Logistic Regression Equation(2Q)

NOT model free -> based on following equations

$$\log \text{ odds} = \beta_0 + \beta_1 X_1 + \dots + \beta_q X_q$$

$$\log p/(1-p) = \beta_0 + \beta_1 X_1 + \dots + \beta_q X_q$$

$$P = 1/(1 + \exp(-(\beta_0 + \beta_1 X_1 + \dots + \beta_q X_q)))$$

EX. Probability could look like this: $P = 1/(1 + \exp(0.6535 - 0.4535(X)))$ ----> where you can sub in x

Direct interpretation of beta 1 is that per unit increase of X1, log odds will increase by beta 1 -> not clear so thus you must say

The Log odds are going to increase by beta 1

Types of Regression :

- Logistic regression (Binary outcome)
- Multiple Linear Regression (Continuous outcome)
- Multinomial Logistic Regression (categorical outcome of 3 or more levels)
- Ordinal Logistic Regression (Categorical outcome of 3 or more ordinal levels)
- Poisson Regression (Count outcome)
- Negative Binomial Regression (Count outcome)



By angelica9373

cheatography.com/angelica9373/

Published 19th October, 2024.
Last updated 22nd October, 2024.
Page 2 of 3.

Sponsored by **CrosswordCheats.com**
Learn to solve cryptic crosswords!
<http://crosswordcheats.com>

Chapter 10 (cont)

All those regression models can be called **generalized linear models (GLM)** !

- 3 equations equivalent to each other/
- $Y = 0$ in MLR is never true if Y is binary and thus cannot use this mode
- Since Y is continuous, change Y into P (probability) and it eliminates the error term since you add some randomness

The Odds

Odds of ration is the exponential form of beta

- Beta is your coefficient number on your regression model

The Odds: $p / (1 - p) \rightarrow p$ is probability

The Odds: $e^{(\beta_1)} = 1$

Logit: Log (odds). It takes the values from - infinity to positive infinity. Dependent var.
Probability: $P = (\text{odds}) / (1 + \text{Odds})$

Comparing 2 Models

First criteria, pick the model with the lowest validation error

Second criterion, when the validation errors are comparable, pick the one with few variables

E.g suppose models 1 and 2 have a validation errors 26.2% and 26.3%. Their model sizes are 10 and, respectively. Which model is better?
- Initially go based of lowest validation error but when its too similar (23% and 26% -> its comparable) and thus you go based off of LOWEST model Size

Performance Evaluation

(1) Partition the data into Training and Validation Sets

Training set Used to grow tree

Validation set used to assess classification performance

(2) More than 2 classes ($M > 2$)

Same structure except that the terminal nodes would take one of the m-class labels

Example Problem

Calculation in logistic regression (Two questions)

• Eg1. Given $\text{Prob}(Y = 1) = 0.5$ for an observation, What is the odds?

• Odds = $p / (1 - p) = 0.5 / (1 - 0.5) = 1$

• Eg2. Given a logistic model as follows. What is the probability of Y=1 for an observation of age 30 and experience 10?

Input variables	Coefficient	Std. Error	p-value	Odds
Constant term	-13.20165825	2.46772742	0.00000009	
Age	-0.04453737	0.09096102	0.62439483	0.95643085
Experience	0.05657294	0.09096305	0.5209661	1.05820346

$X0 = -13.20165825 - 0.04453737(30) + 0.05657294(10) = -13.97295$
 $\text{P}(Y = 1) = 1 / (1 + \exp(-X0)) = 1 / (1 + \exp(13.97295)) = 8.55097e-07$

Check if Binary variable or continuous.

If binary: use the odds ratio DIRECTLY and use ____ "times"

If continuous: odds ratio - 1 then convert to %.



By angelica9373

cheatography.com/angelica9373/

Published 19th October, 2024.

Last updated 22nd October, 2024.

Page 3 of 3.

Sponsored by **CrosswordCheats.com**

Learn to solve cryptic crosswords!

<http://crosswordcheats.com>