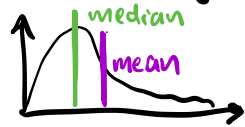


UGBA 104 FINAL

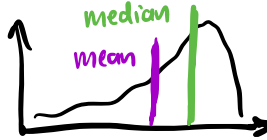
Frequency Distributions

Variables		
Ex:	Customer #	Purchase Item
	1	A
unique observations	2	B
	⋮	15

Histogram Skewness (hint: L/R skew is the opposite of what they look like intuitively).



right skew



left skew

Binning Guide for Quantitative Data

- generally 15-20 bins
- needs to be ME & CE
- Apprx bin width = $\frac{\text{Largest Val} - \text{Smallest Val}}{\# \text{ des. bins}}$

CROSS-TABULATION

Normal Distribution:

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{x-M}{\sigma}\right)^2}$$

σ : standard dev.

M : mean/expected value

Standardization of Data:

$$x_{\text{std}} = \frac{x - M}{\sigma}$$

AKA: "z-score": the # of stan. devs away from the mean a datapoint is.

"outlier" $\hat{=}$ Any datapoint with a z-score outside of $[-3, 3]$

Essence of cross-Tab: slice & dice data into segments & analyze individually

Simpson's paradox: Trends in the "sliced & diced" data may disappear in whole dataset

CLUSTERING

Euclidean Distance: how we measure "similarity" between observations:

$$d_{uv} = \sqrt{\sum_{i=1}^n (u_i - v_i)^2}$$

$$d_{uv, \text{stan}} = \sqrt{\sum_{i=1}^n (z_{ui} - z_{vi})^2}$$

standardized cluv is better otherwise scale of raw data is misleading.

K-means Clustering

INPUT: # of clusters $K \rightarrow$ OUTPUT: cluster assignments.

- ① Assign clusters randomly
- ② Find centroids of each cluster
- ③ Reassign to closest centroid
- ④ Re-calculate centroid location
- ⑤ Repeat until state convergence

Notes:

- converges by definition
- does not always reach global optimum
- "UNSUPERVISED" ML

Business context consid.

- what variables are relevant for the business question?
- what clusters make sense?
- what clusters align w/ expert domain knowledge?

ASSOCIATION RULES

Rule $\hat{=}$ If {Antecedent} then {consequent}

Support Count $\hat{=}$ # of transactions where {A, C} occurs = count(A & C)

Support $\hat{=}$ % of transactions that include item set {A, C} "X % of data exhibits A, & C"

$$P(A \& C) = \frac{\text{support count}}{\text{tot. count.}}$$

Confidence $\hat{=}$ $P(C|A) = \frac{P(C \& A)}{P(A)}$ = support of rule / support of A "X% of data that has A also has C"

// can be misleading if consequent C is too frequent.

$$\text{Lift Ratio} \hat{=} \frac{P(C|A)}{P(C)} = \frac{P(C \& A)}{P(C) * P(A)} = \frac{\text{confidence of rule}}{\text{support of consequent}}$$

Lift Ratio
 $= 1$
 > 1
 < 1

Interpretation
A & C are statistically independent
"Positive Association between A & C" $\rightarrow$ Rule is better than guess
"Negative Assoc. btw. A & C" $\rightarrow$ Rule is worse than guess

Properties

- ① Comparing 2 rules, it's poss. to have $\uparrow$ confidence but a $\downarrow$ lift ratio
- ② if $\{A\}$ then $\{B\}$ & if $\{B\}$ then $\{A\}$ will have the SAME lift ratio

What Rules should we go after

- ① Common ($\uparrow$ support)
- ② Reliable ($\uparrow$ confidence)
- ③ Better than guess (Lift Ratio > 1)
- ④ Actionable
- ⑤ Non-obvious

BINARY CLASSIFICATION

Input Variables $\rightarrow$ Output Variable

Classification vs Regression.

Confusion Matrix

		Predicted Class	
		0	1
Actual Class	0	TN (n_{00})	FP (n_{01})
	1	FN (n_{10})	TP (n_{11})

Performance Measures

$$\text{Accuracy} = \frac{n_{00} + n_{11}}{n_{00} + n_{10} + n_{01} + n_{11}}$$

$$\text{Overall Error Rate} = \frac{n_{01} + n_{10}}{n_{00} + n_{10} + n_{01} + n_{11}}$$

$$\text{Sensitivity} \Leftrightarrow \text{Recall} = \frac{n_{11}}{n_{10} + n_{11}}$$

$$\text{Specificity} = \frac{n_{00}}{n_{01} + n_{00}}$$

$$\text{Precision} = n_{11} / (n_{11} + n_{01})$$

Notes

- false positive $\neq$ false negative, need to look at context
- Class Imbalance: too FEW "positive" examples to learn from
- DO NOT EVALUATE Performance on TRAINING SET.
- Classification is SUPERVISED ML

LOG-REG

Basis LogReg function = $(1 + e^{-(b_0 + b_1x_1 + \dots + b_nx_n)})^{-1} = \hat{p}$

Output: probability $\hat{p}$ ($0 < \hat{p} < 1$) $\approx$ Inputs: $(x_1, \dots, x_n)$ $\approx$ Tune coefficients: $(b_0, b_1, \dots, b_n)$

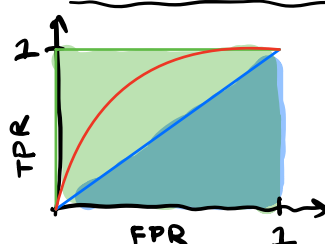
Cutoff probabilities if $\hat{p} > 0.5 \rightarrow$ classify as 1 (0.5 default)

- As cutoff $\downarrow$
- # FP & # TP both increase
- TP rate & FP rate both incr.
- TPR $\Leftrightarrow$ Sensitivity
- FPR $\Leftrightarrow$ (1 - Specificity)

PROS: interpretable
Give association insights to good predictor vars

CONS: relies on a very specific form, fails to capture complex patterns

ROC: Receiver Operating Characteristics



- Perfect classifier
- Random classifier
- ROC curve

The larger $\int$ ROC, the better the classifier

$\int$ ROC = 0.5

$\int$ ROC = 1

k-NN

if $k = 1$: too complex, high variance, overfitting

if $k = N$ = # of samples: too general, high bias, underfitting

PRACTICE

(79, 95) (129, 35) (88, 148)

A

B

C

which pair most similar?

mean var 1: $\frac{79 + 129 + 88}{3} = 98.667$

mean var 2: 92.6

std var 1: = 26.65 // MATLAB

std var 2: 56.53

$$d_{AB} = \sqrt{(1.1381 + 0.7379)^2 + (1.02 + 0.0413)^2} = 2.1554$$

$d_{AC} = 0.9964 \rightarrow$ shortest dist. A & C are most similar (Ellen & Taylor Swift)

$d_{BC} = 2.522$

standardized data

A: (-0.7379, 0.0413)

B: (1.1381, -1.02)

C: (-0.4002, 0.9787)

*NOTE!: you don't have to go through standardization if both variables are absolutely most similar! (saves a lot of work)