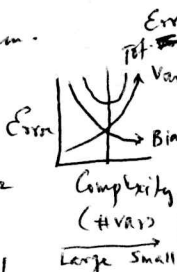


VAR. SELECTION

more $X \Rightarrow \downarrow \text{Bias} \uparrow \text{Var}$
 $\rightarrow$ High Dimensionality
 $\rightarrow$ multicollin
 $\rightarrow$ Pred (S) Explainer
 too few $X \Rightarrow \uparrow \text{Bias} \downarrow \text{Var}$
 too few $X \Rightarrow \uparrow \text{Bias} \uparrow \text{Var}$
 Don't include vars for pred! Want small models!

- # kind of "Data mining". Unresolved problem.
- Nothing special about related vars
- Bias underfitting, Var: overfitting (future resp)



Measure Pred Risk $R(s) = \frac{1}{n} \sum [\hat{Y}_i(s) - Y_i^*]^2$

Don't have future, so $R_H(s) = \frac{1}{n} \sum [\hat{Y}_i(s) - Y_i]^2$ use again! Sneoping!

No penalize: $R_H(s) + \text{Penalty}$ Complexity (C_p) Higher is worse!

Complex Model $\Rightarrow$ down $\uparrow \text{Var}$
 Mallows $C_p = \frac{2/S/\hat{\sigma}^2}{n}$ # of sub-model preds Est var of FULL MODEL ONLY IF NO CONTROL VARS!!

AIC = $2/S/\hat{\sigma}^2$ True var (replac with full or sub-model var)

BIC = $1/S/\hat{\sigma}^2 (\log(n))$ Same as AIC, penalize by $\log(n)$ PREFERRED, PENALIZE LARGE MODELS

LOOCV $R_{CV}(s) \approx R_H(s) + \frac{2/S/\hat{\sigma}^2(s)}{n}$ SUB-MODEL! penalize less than Mallows!

For GLM, use AIC, BIC ("function of tech-valued function")

USE Heuristics if p is too large: Regularized Regression (2^p submodels)

Greedy: SELE on biggest & smallest jump in sel. process
 FOR $\neq$ BACK (maybe), not optimal model
 PREFER! Smaller models BACK: Expensive, no-go if $p > n$ larger model for larger p

$R(s) = \text{Bias}^2[\hat{Y}(s)] + V[Y^*] + V[\hat{Y}(s)]$ Overfitting
 underfitting irreducible! Don't know!

MSE! Can find model w/ lower MSE than full model adding bias reduces MSE! often...

Minimize OLS + λ (penalty) λ balances Lack of fit Model complexity

penalty: L_0 : # of non-zero β s Search all submodels
 L_1 (LASSO) $\|\beta\|_1 = \sum |\beta_j|$ Rewards sparsity Shifts on Multicollin Forces zeroes NO CLOSED FORM!

L_2 (RIDGE) $\|\beta\|_2 = \sqrt{\sum \beta_j^2}$ Deal with Multic Deal with multicomp. easy than Lasso CLOSED FORM

SHRINKAGE VAR SELECTION
 LASSO, RIDGE LASSO (can't get 0) (can't get 0)

- λ determined by k-fold CV (LOOCV is $k=n$) Elastic Net is Both: Ridge dominant if multicoll
- MUST scale predictors! Then interpret intercept of originals
- Less efficient than OLS (increase bias, decrease var.)
- MUST redo OLS with results!
- Evaluate selection MSE (normal Reg) Classification Err (logistic) S. of Sq. Deviances (poisson)
- L_1 : Sparse Model L_2 : Remove limit on # of vars Encourage group effect Stabilizes L_1 path

LOGISTIC REGRESSION ASSUME: 1 Linear 2 Independent

Linear combi of preds $E[Y|X] = p_0 + p_1 X_1 + \dots + p_k X_k$
 Response vars $Y_1, \dots, Y_n$ are indep. Random Vars
 NO ERROR TERM! PROBABILITY OF SUCCESS! NOTHING TO DO W/ RESP!
 $\ln\left(\frac{p}{1-p}\right) = \beta_0 + \beta_1 X_1 + \dots + \beta_k X_k$ $p = \frac{e^{\beta_0 + \beta_1 X_1 + \dots + \beta_k X_k}}{1 + e^{\beta_0 + \beta_1 X_1 + \dots + \beta_k X_k}}$
 Log odds e^{β_j} : β_j unit increase, LOG ODDS RATIO!!

parameters to estimate are $\beta_0, \dots, \beta_k$ (no variance!)
 Done using MLE - NO CLOSED FORM! - log(likelihood) is non-linear!
 Good for LARGE sample sizes, else Type I Error!
 Sampling dist of β is $\sim N(\beta, V)$ no closed form, need software
 HYP. TEST: Use Wald test (z-test really: $n \rightarrow \infty$, large samples) no t-test!

$P(\text{Type I Error}) > \text{Sign. level}$ for small samples...
 ALL APPROXIMATE!! NOT GOF TESTS!

Deviance = $\log(\mathcal{L}_{\text{reduced}}) - \log(\mathcal{L}_{\text{full}})$ χ^2 def $\beta_0, \dots, \beta_k$ model
 Overall reg. test H_0 : Adding anything has no effect

Sub set test H_0 : Adding "q" things has no effect
 $\chi^2_{q, \text{df}}$

LOGIT in Classification: Classifier is $h(x_1, \dots, x_p) = \begin{cases} 1 & \text{if } \hat{p} > 0.5 \\ 0 & \text{else} \end{cases}$

Classif. Error = $1 - \text{Pr}(Y=h)$ DOWNWARD BIAS: Sneoping! (Risk is upward)

Training Logit w/ Cross Valid Random subsample (Expensive, no new k-fold LOOCV)
 $\downarrow$ Bias but $\uparrow$ Var

GOF Testing: All about Residual Analysis. Can ONLY do w/ Binary data w/ Replications! (Bernoulli vs Binomial dist)

Residuals: 1 Pearson: $r_i = \frac{Y_i - \hat{p}_i}{\sqrt{\hat{p}_i(1-\hat{p}_i)}}$

2 Deviances! $d_i = \text{Sign}(Y_i - \hat{Y}_i) \sqrt{\dots}$ Y_i is "saturated model" just observations! $\hat{Y}_i$ is fitted model.

Both Pearson and Dev. have $\sim N(0,1)$ St. Norm Dist iff Assumptions hold! (Check w/ hist., qq plot)

HYP TEST: H_0 : Model fits Data WANT BIG p-value! ONLY TIME we want this!

test stat in D = $\frac{\chi^2}{df} \sim \chi^2_{n-p-1}$

Subset testing is about PRED POWER - Model is good at this even if does not meet assumptions!

GOF is all about testing

Fixing Problem: 1 Add Predictors 2 Transform vars 3 Outliers? (4) Maybe not appropriate for Binomial! Corr among responses

OVER DISPERSION is caused. It means that Heterog in Successes has not been modeled

"The variability of the prob estimator is larger than would be implied by a Binomial Rand. Var"

To fix, maybe another s-function: 1 Probit: Inverse of N-D CDF Tight shape the... 2 Complementary log-log (Good for skewed! long tails!)

But logit is still best. It is "Canonical" $\Rightarrow$ The $\hat{\beta}$ estimates are "fully efficient". Easy to interpret coeffs in terms of log-odds, cannot do with others.

GLM: $g[E(Y|X)] = \beta_0 + \beta_1 \dots \beta_p X_p$

Y_n is a Random Var from an Exp. Dist. family

Model some g as the Expectation of Y as a linear combo of X : $E(Y|X) = \hat{g}(\text{---})$

g is link
 $\hat{g}$ is inv. line

DIST	LINK	REG
Norm	$g(m) = m$	$m = X^T \beta$
Poisson	$g(m) = \log(m)$	$m = e^{X^T \beta}$
Bern/Bino	$g(m) = \logit(m)$ $(\frac{m}{1-m}) \log$	$m = \frac{e^{X^T \beta}}{1 + e^{X^T \beta}}$
λ	$g(m) = \frac{1}{m}$	$m = \frac{1}{e^{X^T \beta}}$

$E = V$ for Poisson...

PDF: $P(Y=y) = \frac{e^{-\lambda} \lambda^y}{y!}$

$E(Y) = V(Y) = \lambda$

Model: $\log(\lambda) = \beta_0 + \dots + \beta_p X_p$

"log rate" $\rightarrow$ When interpreting, RATIO OF RESP RATES

- MLE, no closed form, all approxes...
- Interpret e^β : $\forall$ unit inc in β , the log ratio increases by β .
- Eg: "log of Expected blah" will increase by β
- "The rate ratio of blah" will inc. by e^β

MULTICOLL Linear dep. in design matrix X . Prob. is that $X^T X$ must be invertible for OLS to work. MC prevents this $\Rightarrow$ SE of coeffs and SE of Reg are ∞ (theoretically)

$\rightarrow R^2$ DOES NOT CHANGE! LOL! $\rightarrow \hat{\beta}$ is unstable w.r.t to resp.

$\rightarrow$ SE is artificially large $\rightarrow \hat{\beta}$ has wider CI's

$\rightarrow \hat{\beta}$ less likely to be eval'd as stat. sig $\rightarrow$ F-stat is big but t-stats are not.

1 Pearson's coeff or 2 VIF: "Prop. increase in $Var(\hat{\beta})$ to if it were uncorrelated" (to compare β 's)

$VIF_j = \frac{1}{1 - R_j^2}$ } Just Regressing against other $\hat{\beta}$'s

OK if $VIF_j < \max(10, \frac{1}{1 - R^2_{model}})$ or 1 (lowest)

MLR, t-stat = $\frac{\hat{\beta} - c}{\sqrt{\hat{\sigma}^2 / S_{xx}}}$

Largt t-stat $\Rightarrow$ Smaller p-value!

Prediction outside model scope (new setting) is EXTRAPOLATION.

Prediction } MLR equals

Modeling }

Hypo testing } (Assessment)

ANOVA

Within

Between

$\frac{R^2 - 1}{N - R}$ } for F-stat

$\frac{N - R}{N - 1}$

HOLDING ALL OTHERS FIXED!

Logistic, Poisson: "Null Deviance": $n - 1$

"Residual Deviance": $p - p - 1$

Wald test: $Z\text{-stat} = \frac{\hat{\beta}}{SE(\hat{\beta})}$

Reject H_0 if $|Z\text{-stat}|$ is too large $\Rightarrow \hat{\beta}$ is sig.

$H_0: \beta \leq 0$ } t.e. for

$H_A: \beta > 0$ } sig +ve

p-val = $P(Z > Z\text{-val})$

"Explained by Regress"

$R^2 = \frac{SSR}{SST}$

ANOVA est. simultaneous conf Int.

R^2 is Coeff of Determination = $1 - \frac{SSE}{SST} = \frac{SSR}{SST}$

F-stat = $\frac{MSR}{MSE} = \frac{SSR/p}{SSE/(n-p-1)}$, Adj $R^2 = [1 - (R^2)(\frac{n-1}{n-p-1})]$

Cooks measure how much param change when 1th obs. is removed: $\frac{4}{n}$ } if > than this, investigate

F-stat = $\frac{\text{Bet. Grp}}{\text{Within Group}} = \frac{\text{Explained Var}}{\text{UnExplained Var}} = \frac{SSR/(k-1)}{SSE/(n-k)}$ } FOR ANOVA!

Causality $\rightarrow$ ISOLATION. ASSOCIATION is no ISOLATION

ASSUMPTIONS

Simple Linear: linearity (mean zero) $E(\epsilon_i) = 0$

const. var $Var(\epsilon_i) = \sigma^2$

Indep. ϵ_i is ind

Normality: $\epsilon_i \sim N(0, \sigma^2)$

SCATTERPLOT

RES ANALYSIS

Linearity

Cons var

Indep

Cons Var

Indep (Uncorrelated NOT independent errors!)

MLR: Std Res vs Var $\rightarrow$ LIN

SH Res vs FIT $\rightarrow$ CV, IND

QQ, Hist $\rightarrow$ NORM

$\lambda = 1$ none

$\lambda = 0$ $\log(Y)$

$\lambda = -1$ $1/Y$

$\lambda = 1/2$ $\sqrt{Y}$

REPAIR: Norm } Transform Y

Cons var } (next) Box Cox

Linearity } Transform X

Simple: one QUANT

ANOVA: one QUANT

MLR: Every thing

TRUE

$E(\epsilon_i) = 0$

$Var(\epsilon_i) = \sigma^2$

HARRY FORM

$E(\epsilon) = 0$

$Var(\epsilon) = \sigma^2 I$

$E(\hat{\epsilon}) = 0$

$Var(\hat{\epsilon}) = \sigma^2 (I - H)$

USE RESID NOT RESP

EST

$E(\hat{\epsilon}_i) = 0$

$Var(\hat{\epsilon}_i) = \sigma^2 (1 - h_{ii})$

ANOVA: pooled var est has $(N-k)$

MLR: $\hat{\sigma}^2 \propto MSE = \frac{\hat{\epsilon}^T \hat{\epsilon}}{n-p-1} \sim \chi^2_{n-p-1}$

MAE = $(1/n) \sum |Y - \hat{Y}|$

MSE = $(1/n) \sum (Y - \hat{Y})^2$ less robust to outliers

MAPE = $(\frac{100}{n}) \sum \frac{|Y - \hat{Y}|}{Y}$

Precision Measure = $\frac{\sum (Y - \hat{Y})^2}{\sum (Y - \bar{Y})^2}$

Est. Residuals in MLR do not have cons. var!

(Divide by Stddev and standardize)

$\hat{\epsilon}_i = \frac{\hat{\epsilon}_i}{\sqrt{\sigma^2 (1 - h_{ii})}}$

(correlation coeff) = $\rho = \sqrt{R^2}$

Variable Selection

- ① **Avoid overfitting**: random effects are given a chance to rear their heads
- ② **Simpler models**: Easy to explain and less chance of including insignificant factors
- ③ **Can be illegal to use certain factors**: Remember to remove other factors that correlate strongly (eg. UT and Memory)

How: Greedy Algos

Consider only one thing at a time and do not consider future options.

① Forward Selection, ② Backward Elimination, ③ Stepwise regression

Start w/ some or all factors → add/remove gradually

remove factors ← fit model ← "Is factor good enough?"

w/ $p > 0.05$ (p-value) or $p < 0.05$ (p-value)

Filters: p-value, R^2 , AIC, BIC

Global Algos: * Need to scale data! This is Optimization!

① LASSO: $\sum_{i=1}^n |a_i| \leq T$ → "budget" pick many diff. values

② Elastic Net: $\lambda \sum_{i=1}^n |a_i| + (1-\lambda) \sum_{i=1}^n a_i^2 \leq T$

③ Ridge Reg: $\sum_{i=1}^n a_i^2 \leq T$ (does not do var. selection!)

Comparison: Greedy are good for initial analysis but fit random effects more

Lasso: Good for var. selection. $\sqrt{\frac{1}{n}} \frac{1}{\sqrt{1-\alpha}} \sqrt{\frac{1}{n}} \frac{1}{\sqrt{1-\alpha}}$ will just pick one

Ridge: Good for prediction $\square$ ← will push down, so $\square$ power

Pred. error = f(bias, variance) ← Tradeoff bias for variance for better model.

DOE: Data is not/never available. Compare two factors by controlling for blocking factor that creates variance.

① A/B Testing: Compare two alternatives. * Run hyp. tests after higher order possible only datapoint collection!

② Factorial Design: Compare combos. Full and Fractional (subset of full)

Balanced Design: a) Each choice same # of times b) Each pair of choices same # of times

- Independent factors? Use fractional w/ Regression.

- Exploration ("gather data") vs Exploitation ("use it")

"How quickly are you able to decline value?"

③ Multi-Armed Bandit: Compare alternatives. Do both explor and Exploit. Old ladies in Vegas.

→ Alternatives: A/B, High order A/B, C/D, E/F, Full, Fractional

PROB-BASED MODELS

① Bernoulli: "Will x happen or not?"

PMF: $P(X=x) = p^x (1-p)^{1-x}$ * For large n → Normal Dist

iid! Data cannot "clump" eg. temp over a year

② Binomial: "x successes out of n Bernoulli trials"

PMF: $P(X=x) = \binom{n}{x} p^x (1-p)^{n-x}$ * For large n → Normal Dist

③ Geometric: "How many trials until..."

PMF: $P(X=x) = (1-p)^{x-1} p$ * Careful defining "success"

* Check iid! If Geometric, is iid

④ Poisson: "x things happening w/ AVG happen rate λ "

PMF: $P(X=x) = \frac{\lambda^x e^{-\lambda}}{x!}$

⑤ Exponential: "If λ follows dist, it is Poisson"

PMF: $P(X=x) = \lambda e^{-\lambda x}$ "If Poisson, λ is exp. dist"

⑥ Weibull: "Time until something happens"

AND "Time bet. failures" PMF: $P(X=x) = \frac{k}{\lambda} \left(\frac{x}{\lambda}\right)^{k-1} e^{-\left(\frac{x}{\lambda}\right)^k}$

$k < 1 \Rightarrow$ "Worst things fail first" ~~Failure rate DEC w/ time~~

$k > 1 \Rightarrow$ opposite eg. tires, bulbs. Failure rate INC w/ time

$k = 1 \Rightarrow$ fail rate constant $\Rightarrow$ Exponential.

• Geometric: "# of flips until bulb goes out"

• Weibull: "amt of time until bulb goes out"

Q-Q Plot: ① Compare two dataset percentiles ② Map dataset to a prob. dist. why? a) Tests don't give visual info b) We might be interested in just the extreme values!

Queueing: Needs something resembling a queue.

Memoryless: Next state only depends on current. Follows Exp, Poisson.

if not, use Weibull (eg. bulb, tires)

Queueing theory, Kendall notation is standard

Simulation: When things don't follow prob. dist.

→ Deterministic Continuous time (PDEs are used)

→ Stochastic Discrete event Stochastic Sims (DESS)

DESS: ① Build model w/ a) Entities (people, bags) b) Modules (queue, storage) c) Actions d) Resources (waiters, grinders)

② Decision points ③ Replication: run many times to get dist

③ Validate w/ real data * Careful can have same or diff. σ^2

④ Use: Simulate w/ diff options. Use same randomness!

Markov-Chain Models: * Assume system is memoryless! (Based on system status) * No cycles! * Every state reachable from other.

$\pi \rightarrow$ state vector

$P = \{P_{ij}\} \rightarrow$ Trans. Matrix w/ state trans prob P_{ij} . Solve for

$\pi^* S: \pi^* P = \pi^* \leftarrow$ "Steady State" ← trans. don't change

Eigen vector w/ Eigenvalue of 1. Eg. Page Rank by Google

Missing Data & Imputation: - Data wrong or missing - Some data more than others (bias!)

Non-Estimation methods: ① Throw away (not putting in error)

② Cat. vars w/ interaction terms. (like getting two models) (also doubling vars)

Estimation Method: ① Simple: Data numeric: Mean, Median Catey: Mode

② Predictive: simpler but bias possible to under/over estimate

eg. w/ regression - Not capturing variability

③ Perturbation: Add rand. to each imputed val from some dist.

DISTURB: AVG, VAR | DONT DISTURB AVG, VAR

* Don't impute more than 5% of data, many factors

Error: (Impute Error) + (Perturb Error) + (Model error)

But all data has error anyway!

ADV MODELS: What if no good models? Hyp. tests assume this but what if none?

Unknown dist? Non-param tests: - Not much data - Need median - Unknown dist

① McNemar's Test: Match data in resp. (paired, matched) Throw away stuff that matches Use binomial to compare discordant.

② Wilcoxon Signed Rank: Assume response is cont., symmetric (single, paired, matched) "Is median different than m?" means below III they are w/o distribution!

③ Mann-Whitney Test (unpaired) ④ Bayesian Models: Very little data Est. distribution.

$P(A|B) = \frac{P(B|A) \cdot P(A)}{P(B)}$

Combine w/ expert opinion to get Single observ. can be related to larger!

$P(A) =$ "Prior dist"

$P(A|B) =$ "Posterior" (new dist)

$P(A^c) \cdot P(B|A^c)$

Network Models: "Find highly interconnected subpopulation"

Louvain Algo: Heuristic. Maximize "Modularity"

$$Modularity = \frac{1}{2W} \sum_{i,j \text{ in same community}} (a_{ij} - \frac{W_i W_j}{W})$$
 Weight of edge (can be zero)
 Weight of edges to node
 Wt of all edges

If trying to minimize each node stays its own community

Neural Nets: Input $\rightarrow$ Hidden $\rightarrow$ Output "Neuron"
 Loop, keep adj. weights (eg. w/ Gradient descent), pick output with highest result.

They suck $\leftarrow$ Lots of data req. to train
 Learning algo tough to tune so weights don't change slowly

Competitive Decision Making: Game Theory! Eg. Adjust "Co-operative Game Theory" model and competition in compete + co-operate. reacts.

- Strategies $\leftarrow$ Pure: only one choice
 Mixed/Random: Rock, paper, scissors

- Information $\leftarrow$ Perfect: know all about opp.
 Imperfect: Don't

- Sum $\leftarrow$ Zero: "Benefit" = 0. One win $\Rightarrow$ One loss
 Non-zero: "Benefit" > 0 or < 0 .

Optimization models used.

OPTIMIZATION DYNAMIC PROG. $\leftarrow$ States/Decisions

- Basic Dyn. Prog (State/decision) $\leftarrow$ Optimization thinks that all data is known.
 $\downarrow$ Add prob. of state transitions
- Stochastic Dynamic Program. $\leftarrow$ Bellman's Eq to optimize
 $\downarrow$ w/ discrete state and decision
- Markov Decision Process (Prob. depends only on current state & decision)

MATHEMATICAL MODEL PROGRAMS.

* We still need models! Only solver stuff!

$\rightarrow$ Variables: Decisions to make
 $\rightarrow$ Constraints on variables

$\rightarrow$ Objective Function: Measure quality of soln.
 * Can incl. inputs from other models!

$\rightarrow$ Solution $\leftarrow$ "Feasible" if satisfies constraints
 "Optimal" if "Feasible" and best Obj. Val.

Binary Vars: Create complex models.

① Define in VARS, ② Add to OBJ FUN, ③ Link in constraints

Example: Max. of any food is M. $x_{cs} \leq M y_{cs}$ (if eat then one or the other $y_a = 1 - y_b$)
 If one then other is a must $x_a \leq M y_b$
 both or none $y_a = y_b$
 or: $y_a + y_b = 1$

Model Classification: Linear $\rightarrow$ Conv Quad $\rightarrow$ Conv $\rightarrow$ Int (Eg, quick) $\rightarrow$ Gen. Co. (hard)

Linear: $f(x)$ and CONS Linear $\sum a_{ij} x_i \leq b_i$ ($>, =$)

Convex Quad: $f(x)$ convex, CONS linear. $\sum c_i x_i^2, \sum a_{ij} x_i \leq b$

Convex: $f(x)$ convex concave, CONS convex $\sum c_i |x - z|, \sum a_{ij} x_i \leq b$

Integer: $f(x)$ linear but CONS has all or few integers. if $x \in \mathbb{Z}$ or $\{1, 2, \dots\}$

General Non-Convex: Just not convex eg. $\sum c_i \sin^2 x_i$

Network Model: flow through node } linear function
 flow through arc } in OBJ FUN

Shortest Path
 Assignment Model (workers, eq)
 Max Flow model

Uncertainty: ① Model conservatively (eg. add 0)
 ② Scenario modeling d_i demand on day i scenario

Statistical Model: variable in x_{ij}
 coeff in a_j

Optimization Model: variable in a_j
 coeff in x_{ij}

Lasso: $\sum |a_i| \leq t$ **Ridge**: $\sum a_i^2 \leq t$ **Elastic Net**: $\lambda \sum a_i + (1-\lambda) \sum a_i^2 \leq t$

LOGISTIC: VARS $a_0 \dots a_m$ CONS none OBJ FUN

$$\text{MAX} [\prod p(x_i) \prod (1-p(x_i))]$$

$$p(x_i) = \frac{1}{1 + e^{-x_i}}$$
 (min. pred. error.)

SVM Hard: $a_0 \dots a_m$ $(A_0 + \sum a_j x_j) y_i \geq 1$ $\text{MAX} [\sum a_j^2]$
 (min margin/dist bet. vectors)

SVM Soft: $A_0 \dots A_m$ none $\sum \max \{0, 1 - (A_0 + \sum a_j y_j)\} + \lambda \sum a_i^2$

Exp & Smooth: α, β, γ $0 \leq \alpha \leq 1, 0 \leq \beta \leq 1, 0 \leq \gamma \leq 1$ $\sum_{t=1}^n (y_t - \hat{x}_t)^2$
 (min pred. error)

ARIMA Time Ser.: μ, ϕ, θ none

GARCH: ω, β, γ_i none $\sum_{t=1}^n (\sigma_t^2 - \hat{\sigma}_t^2)^2$

CLUSTER K-means: z_{ik} cluster ctr y_{ik} point $\sum y_{ik} = 1$ (all assigned to cluster)

$$\text{MM} \sum_i \sum_k y_{ik} \sqrt{\sum_j (x_{ij} - z_{jk})^2}$$
 (min. dist to cluster ctrs)

⑥ R^2 is frac. of variance explained by model.

⑥ Overfitting on training $\Rightarrow$ Bad perf on Val and Test
with meaning - independent

② Data (structured) $\left\{ \begin{array}{l} \text{# with meaning} \\ \text{quantitative} \\ \text{# w/o meaning} \\ \text{categorical} \end{array} \right. \left\{ \begin{array}{l} \text{independent} \\ \text{time-series} \end{array} \right.$

Binary $\left\{ \begin{array}{l} \text{text is unstruc.} \\ \text{data} \end{array} \right.$

• **SVM** (Classification): the ~~the~~ $\sum a_i$ is to be min to max the margin  (dist. bet. lines $\propto \sum a_i$)
optimization: $\text{Min} [\text{Max}(0, 1 - \sum (a_i + a_i x_i) y_i) + \lambda \sum a_i]$

Com Split that: $25 \sum_{i=1}^{25} + 20 \sum_{i=26}^{30} + 15 \sum_{i=31}^{47}$ } Want to avoid errors in this range more

- Scaling between $[0, 1]$ to $[b, a]$: $Sx_j = \frac{x_j - \min(x)}{\max(x) - \min(x)} \cdot (a - b) + b$

• Standardize:
$$z_{ij} = \frac{x_{ij} - \mu(x)}{\sigma(x)}$$
 Can be zero!

[k-NN] (Classification, Resp. Pred.) new point y_i
 closeness, $d = \sqrt{\sum (x_i - y_i)^2}$ weighted $d = \sqrt{\sum W_i (x_i - y_i)^2}$

• **K Means** (Clustering) (heuristic, Expectation Max)
optimization prob: "minimize dist. to cluster center iff point is in cluster."
Sum of dist to c_k

Can do predictive clustering
w/ Voronoi diag. 

CUSUM (Change Detection) Done w/ time-series data
Hyp. testing might be slow.

Inc $S_t = \max\{0, S_{t-1} + (x_t - \mu - C)\}$ → pull down sensitivity
Alarm if $S_t \geq T$ (threshold) → high val → more hard to detect → Reduces correlation

$S_t = \text{MAX} \{0, S_{t-1} + (\mu - x_t - c)\}$ $\rightarrow$ less take down
 "Which factors are important?" $\rightarrow$ Remove randomness
 Rank by importance
 Simpler models!
 (use n top PC's)

PCA
 ① Scale Data, ② Normalize mean
 $S \leq \sigma_i / n = \mu_i = 0$

③ Get an Eigenvector with sorted Eigenvalues.

④ Use regression to find new co-effs. for new factors
 ✓ compos of old ones!!

New factors are linear combos of old ones.
 k^{th} factor for i^{th} datapoint $= t_{ik} = \sum_j x_{ij} v_{jk}$
 $y_i = b_i + \sum_{k=1}^K b_k t_{ik}$ ($b_0, b_1, \dots, b_K$) $\rightarrow$ not the same #
 orig factors! (implied reg. coeff) $a_j = \sum_{k=1}^K b_k v_{jk}$ (Do the math) as orig. model
 using new, old values.

like time-series data (ONLY Y)

Outliers:

point

Contextual
(Some things different)

Collective
(many are different)

Reaching Output : p -value : prob. that coefficients are zero
 t -statistic : p -value / std err

R^2 : "How much variability does model account for?" $R^2 = 0.54 \Rightarrow$ "My model acc. for 54% of variance. 46% was randomness or something else..." (usually) more predictions (or less) mean more error.

Model Quality & Fit : MLE : assume have correct data.
assume known variance.

If $n \cdot D$, Error $\sim N(0, \sigma^2)$. Prob. of seeing z_i if y_i is given by density func (Gaussian): DF so $\text{MAX}[\frac{1}{\sigma} \text{DF}]$
 $\Rightarrow \text{MIN}(Z_i - y_i)^2 = \text{MIN}(Z_i - (a_0 + \sum a_i x_i))$ IF HIDDENLY!
 $\rightarrow$ ML value

AIC: $AIC = 2k + 2 \ln(L^x)$ → penalty term.
 for x points $L^x = ML \text{ value}$ $\left\{ \begin{array}{l} \text{Bal. likelihood w/ simple} \\ \text{avoid overfitting.} \end{array} \right.$

$AIC_C = AIC + \frac{2k(k+1)}{n-k-1}$

$k = \# \text{ of parameters}$
 $n = \# \text{ of data points}$

LOW IS GOOD

Model with

Compare $AIC = \exp\left(\frac{AIC_1 - AIC_2}{2}\right) = x\%$ $\Rightarrow$ Model 1 is $x\%$ more likely to be better than Model 2.
Order of terms matters! (higher penalty term $\Rightarrow$ models with more terms are penalized more)

BIC: Bayesian. Has higher penalty for
fewer params than AIC. Use only when more data points
than params. **ORDER MATTERS!** $BIC1 - 2 \times 10^{10}$
Smaller BIC "very likely" better

Splitting Data : Single Model : $6 < \Delta < 10$ "likely"
 $70/30$, $90/10$
 Tr Ts Tr Ts $2 < \Delta < 6$ "Some what likely"
 Compare Model : $50/70$, $Spl + Ts$, Va $0 < \Delta < 2$ "Slightly likely"
 Tr data equality

- Random: may not rep. data equally
 $T \rightarrow \text{val} \rightarrow T \rightarrow T \rightarrow \dots$
 prob. w/ missing data
- Potate: Building eg: MWF in Tr
 (days/week)

T_{Train} : Fit params
 $T_{\text{Validation}}$: Pick model
 T_{Test} : Est. generalizability
 k-fold x-validation

Cross validation: $(Tr + Va) / Ts$: k -fold cross validation
 k sets. Tr on $k-1$, Va on 1

Split $(T+V)$ into K sets, H
Average across all folds, pick best
Train best again on $(T+V)$!!
training vs test sets.

② Diff. random effects in training.
High-perf model in training may get boost from random effects.
Shallow-based models (Reg. SVM) to handle this.

Detrending Data: Factor-data cannot be sensitive to trend, cannot be sensitive to each factor.

reasonability. So run regression ^(response) and remove from ~~the~~ prediction data is like this

Normality Assumption: Not all have equal variance. We use $t(y) = \frac{(y^2 - 1)}{2}$

Heteroscedastic $\Rightarrow$ non-constant variance
Box-Cox to get some normality
Q-Q plot.

Check with a ϕ - ϕ probe $\rightarrow$ more critical, life and death
Salt and Hard $\rightarrow$ that suffering

② SVM: Soft max
don't put classifier in place that performs
best of misclassifying is higher.

Classification Matrixes: "True Negative"
 Not in category

CONFUSION MATRICES

E.g. $\frac{TP}{TP + FP}$ → model is right
→ 9 out of 10 numbers

Sens. Fidelity = $\frac{TP + FN}{TP + FN + FP}$ } frac. of Category members
correctly identified

Specificity = $\frac{TN}{TN+FP}$ } frac. of NON Cat. members
correctly identified.

② SVM soft: OK to allow for small margins to

reduce class. errors in training set.

0.05 (c) A model tells you nothing about our rep.
You need to prep. data... GIG0, 40.

100

Exp Smoothing (Resp. Pred) Use for time-series

data to predict ~~various~~ baseline
response. But like CUSUM but not really.
Est Expect Observ
 $S_t = \alpha x_t + (1-\alpha) S_{t-1}$ @ time t
baseline
 $0 < \alpha < 1$
"observed is real indicator" $\Rightarrow$ new baseline
 $S_t = x_t$
"random luck... keep same" $\Rightarrow$ same, random luck
 $S_t = S_{t-1}$
So, if too random, keep the old (S_{t-1}) with low luck α
if not too random, keep the new w/ high α (observed value)
 $\alpha \rightarrow 0$ for lots of randomness
 $\alpha \rightarrow 1$ for not much randomness.

Trends: Eg: Spikes during Christmas time.
ADDITIVE! $S_t = \alpha x_t + (1-\alpha)(S_{t-1} + T_{t-1})$
 $T_t = \beta(S_t - S_{t-1}) + (1-\beta)T_{t-1}$

Seasonality/Cycles: Inflate/deflate the obs. x_t ,
MULTIPLICATIVE! not (baseline + trend)
"L" cycles/periods
 $S_t = \frac{\alpha x_t}{C_{t-L}} + (1-\alpha)(S_{t-1} + T_{t-1})$ eg: BP measured every hour/day $\Rightarrow L=24$
 $C_t = Y(\frac{x_t}{S_t}) + (1-Y)C_{t-L}$ $C=1$ no change $C=0.9 \Rightarrow$ 10% lower $C=1.1 \Rightarrow$ 10% higher

see the effect of C_t on the observed value x_t
Forecasting: $F_{t+1} = \alpha S_t + (1-\alpha)S_t$ | $F_{t+1} = S_t$ initially
w/ trends: $F_{t+k} = S_t + kT_t$ $k \in \mathbb{N}$
w/ cycles: $F_{t+k} = (S_t + kT_t) C_{(t+1)-L+(k+1)}$
To find α, β, γ , optimization prob: $\text{MIN}(F_t - x_t)^2$
Good for short-term stuff (Resp. pred, Classification)

ARIMA: Again, time-series only! Use when data is not stationary (ie. μ, σ^2 not constant over time)
But diff. can be stationary. d-order diffs.
3-order diff: $\text{diff}(\text{diff}(\text{diff})) \leftarrow$ chained
(P) time periods, to predict (d)-order diffs
Basically auto-regression applied to diffs.
Uses prev. errors as predictors: (P) time-periods (???)
AR(0,0,0) white noise, (0,1,0): Random walk, (p,0,0): basic auto-reg. model (0,0,q): Moving avg model.

ARIMA(0,1,1) = EXPONENTIAL SMOOTHING!
So far, to forecast values.
Need something to forecast variance.
(Variance Est.)
Need variance $\rightarrow$ estimate error.
 $\rightarrow$ make it proxy for something (eg. investment risk)

GARCH: Need variance $\rightarrow$ estimate error.
"How much is forecast w/ 10% than true"
"use variance (sq. diffs)"
"use RAW! (NOT diff. of variance!)"
"Is also autoregressive."

GARCH(p,q)
So far, to forecast values.
Need something to forecast variance.
(Variance Est.)
Need variance $\rightarrow$ estimate error.
 $\rightarrow$ make it proxy for something (eg. investment risk)

GARCH(p,q)
So far, to forecast values.
Need something to forecast variance.
(Variance Est.)
Need variance $\rightarrow$ estimate error.
 $\rightarrow$ make it proxy for something (eg. investment risk)

GARCH(p,q)
So far, to forecast values.
Need something to forecast variance.
(Variance Est.)
Need variance $\rightarrow$ estimate error.
 $\rightarrow$ make it proxy for something (eg. investment risk)

REGRESSION (Classification, resp. pred), (Desc, Pred)

Linear: $y_i = a_0 + \sum a_i x_i$ | $\hat{y}_i$ | Error = $(\hat{y}_i - y_i)$
(pred) (real)
Optimization problem: $\text{MIN}[\sum (\hat{y}_i - y_i)^2] + (a_0, a_1, \dots, a_n), n \in \mathbb{N}$
(Response in Continuous)

Logistic: use for prob. values (can turn to binary classifier eg. $p > 0.5 \Rightarrow 1$)
(Classification)
Linear: $y_i = a_0 + \sum a_i x_i$, take and make $\log(\frac{p}{1-p}) = y_i$
 $\Rightarrow p = \frac{1}{1 + e^{-y_i}}$, when $y_i = -\infty, p = 0$
 $y_i = +\infty, p = 1$
• Pseudo R^2
• Takes longer
• No "cloud form"

Prob. that logit model will give a high P to Someone from YES group (in category) than NO group
No idea of the cost of false +ve, false -ve.
ROC Curve C_i (better!)
Sens (TPRate)
AOC
if = 0.5, is a guess
0 1-Spec (FPRate)
affected by coeffs. / param.

Poisson: if not ND. eg: Count arrivals of people at airport line.
 $f(z) = \frac{\lambda^z e^{-\lambda}}{z!}$ $\rightarrow$ NOT a param! function of something (eg. time)
Est. $\lambda(x)$

Splines: use knots and k-splines (smooth) to fit (non-parametric)
Bayesian: Not much data // use some coeffs, rand. err to start. OR use a "broad prior dist" as seed data. ...

KNN: LOL: Give avg. resp. of all k-neighbors.
(non-parametric) Can weight resp. too.
CART: Split regression models up into tree. Model per leaf. Must have at least 5% of data in each leaf (else overfitting).

in each leaf (else overfitting):
BUILD
① Take half the data: foreach(leaf):
foreach(factor):
compare σ^2
split if $<$ threshold (check size)
ie, pick factor with lowest total variance.
② Take other half.
PRUNE foreach(branch-point):
estimate error (w and w/o branch)
prune if $>$ threshold.
Some est. of quality
"x > 2" prune and make binary id.

Random Forest:
① Bootstrap: Pick from data randomly
Some appear once, some more than once, some not at all
Same size as dataset!
② Branch on random (1+logm) factors. NO PRUNING STEP!
on imp. of each factor based on # of branches.

Regression Tree? Average the responses
Classification Tree? Most common response (mode)
Good: Better overall est. overfitting rbd: of averaging
BAD: Hard to explain/interpret. Cannot point to specific model.