



Centro Stampa

**ATTENZIONE QUESTI APPUNTI SONO OPERA DI STUDENTI , NON SONO STATI VISIONATI
DAL DOCENTE. IL NOME DEL PROFESSORE, SERVE SOLO PER IDENTIFICARE IL CORSO.**

N°1003

**CLASSIFICAZIONE E INTERPRETAZIONE
TEORIA TEMI ESAME**

DI BETTALE VALENTINA

APPLICATION DEVELOPMENT Sviluppo delle Applicazioni

Come possiamo costruire un'applicazione?

Dal problema $\rightarrow$ si arriva al sistema
3 passi

1° ANALISI (concettualizzazione del problema)

Capire e studiare bene il problema x avere le specifiche validate e x sapere quali sono:

- I dati a disposizione
- il ruolo e la gestione dell'interfaccia
- le problematiche e che tipo di approccio bisogna utilizzare

2° Processo di costruzione di un modello

Grazie a matematica, logica, informatica...
 e scelta del metodo con cui utilizzare questo modello

modello e metodo e il tutto deve successivamente essere adeguato al sistema specifico (Tuning)

alla fine si ottiene il prototipo del sistema

⚠ nel caso ci si accorge che manca qualcosa si può tornare indietro alla fase di analisi

3° Processo di verifica e validazione

$\downarrow$

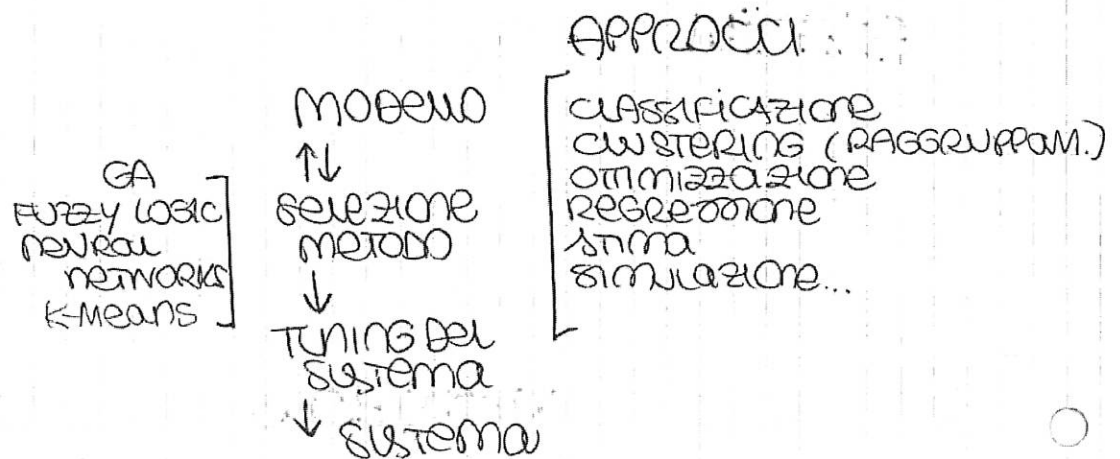
 x il prodotto deve — soddisfare le specifiche (ver.)
 sistema pronto — essere funzionale x chi lo utilizza (val.)

⚠ ogni volta che si è risolto un problema, ce ne si crea un altro (cerchio ∞)

noi partiamo dal punto ②, dal construction process

Processo di Costruzione

- include:
- quale approccio usare/scegliere x risolvere il problema
 - quale metodo e quale modello scegliere
 - il tuning (adeguazione al problema)
 - il sistema ottenuto



1° PROBLEMA

Come si possono descrivere le caratteristiche di un set di elementi?

Come **descrivere** i dati?

Un elemento (o oggetto) è caratterizzato da TNT dati

Si incontra questo problema quando si prova ad analizzare i dati acquisiti durante un TRIAL

(= studio controllato e progettato di elementi selezionati)

La risposta al problema è la **STATISTICA DESCRITTIVA**

Metodi statistici:

Le statistiche possono essere di 2 tipi:

① STATISTICA DESCRITTIVA

Usata semplicemente x descrivere il campione che ci interessa

è usata - innanzitutto x ottenere un primo sguardo sui dati

- poi x usare nel test statistici

- infine x indicare l'errore associato ai risultati e x guaiuti grafici

② Statistica inferenziale

inferenza = deduzione
generalizzazione

Si fanno delle inferenze su come un valore statistico rappresenta tutta la popolazione

serve a valutare la rappresentatività del campione, altrimenti non si può estendere subito un risultato di un trial a tutta la popolazione
→ obbiezione tipica del trial: criteri di selezione x arruolarsi!

Sono tecniche che permettono di usare i campioni x fare delle **generalizzazioni** sulla popolazione
ovvero come sono presi i campioni

⚠ lettere x indicare - popolazione - campioni
 μ, σ x, s, a

[la popolazione include tutti gli oggetti di interesse, mentre il campione è solo una porzione della popolazione.

Parametri → associati alle popolazioni,
lettere greche

Statistiche → associate ai campioni,
lettere romane

Quando si usano i metodi inferenziali, è importante che il campione rappresenti accuratamente la popolazione

5 modi x estrarre un campione (stili di estrazione)

① E. casuale: casualità totale
(RANDOM) ogni elemento nella popolazione ha la stessa probabilità di essere estratto

è il modo preferito di estrazione, ma è molto difficile da fare

② E. sistematica: campionamento sistematico ogni n elementi
è + facile del random
ex: segnavi

③ E. CONVENIENTE : non si cerca un campione specifico, ma chi c'è ogni dato a disposizione è usato
(convenience)

Facile da fare, ma è forse la tecnica peggiore

Semplice da usare anche x l'organizzazione

④ E. CLUSTER : divide la popolazione in gruppi e poi estraggo da ogni gruppo elementi
la popolazione è solamente divisa geograficamente

I gruppi sono chiamati cluster o block

→ si sceglie casualmente un gruppo e si usano tutti i suoi elementi

⑤ E. STRATIFICATO : la popolazione è divisa in gruppi chiamati strati, x qualche caratteristica, non geograficamente
ex: M e F

un campione è preso da ogni strato usando ①, ② o ③

PARAMETRI della STATISTICA DESCRITTIVA

● VALORE MEDIO o MEDIA ARITMETICA ($\bar{x}$)

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$$

Sommativa di tutti gli elementi divisa x il numero di elementi stesso

● MEDIANA

n dispari $\rightarrow \frac{n+1}{2}$

⇒ prendo il valore centrale

n pari $\rightarrow \left(\frac{n}{2} + \left(\frac{n}{2} + 1 \right) \right) / 2$

⇒ prendo la media tra i 2 valori centrali

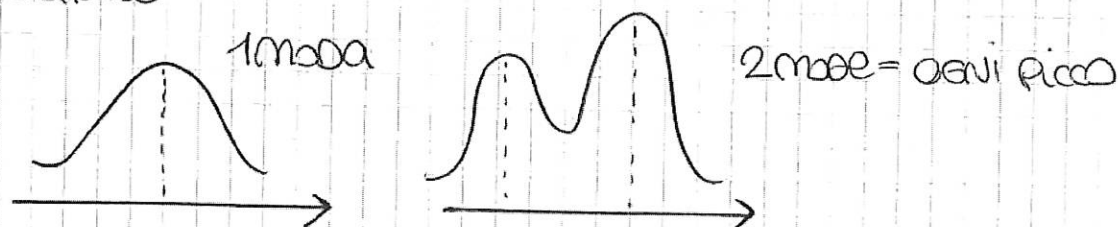
a s_x e a s_x della mediana devo avere lo stesso numero di elementi!

● MODA

valore che si ripete più frequentemente tra tutte le osservazioni in un campione

può non essere unica

valore che + frequentemente assume la variabile



● VARIANZA

dice quanto è dispersa la popolazione:

$S \uparrow$: campione disomogeneo

$S \downarrow$: campione omogeneo $\rightarrow \bar{x}$ significativo

$$S^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$$

Dipende dall'unità di misura della variabile!

● Deviazione Standard

$$S = \sqrt{S^2}$$

● Coefficiente di Variazione

$$100\% \cdot \left(\frac{S}{\bar{x}} \right) \quad \text{valore adimensionale!}$$

→ risolvere il problema della dipendenza e confrontare come valgono 2 valori e la loro riproducibilità

questo valore rimane sempre lo stesso indipendentemente da quali unità di misura si usano

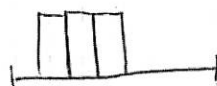
● RANGE

Differenza tra l'osservazione massima e la minima

coppia formata dal valore minimo e dal valore massimo che trova nel campione
da l'idea di quanto è ampio il campo di valori della mia variabile

● PERCENTILI

SI PU' DIVIDERE L'INTERO INTERVALLO DI VALORI IN CLASSI
CON LO STESSO NUMERO DI ELEMENTI
IL PERCENTILE FA SI' CHE LE BARRE SIANO TUTTE ALTE UGUALI,
CIOE' CONTENGANO LO STESSO NUMERO DI ELEMENTI.
Δ NON TUTTI GLI INTERVALLI SARANNO UGUALI, ANZI SONO ESSI
CHE COMBINO X AVERE IN OGNI BARRA LO STESSO NUMERO DI
ELEMENTI



● DISTRIBUZIONE DELLA FREQUENZA

E' una visualizzazione ordinata di ciascun valore
e della sua frequenza, cioè quante volte
appare nel set di dati

PUO' ESSERE ESPRESSO COME UNA PERCENTUALE.

UTILE CON VARIABILI DISCRETE

DATA UNA DISTRIBUZIONE GAUSSIANA SI DEFINISCONO ANCHE:

● SKEWNESS : MISURA DELLA MANCANZA DI SIMMETRIA

MISURA SE LA DISTRIBUZIONE E' SIMMETRICA O NO
RISPETTO AL VALORE MEDIO

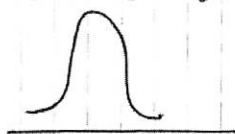
$$SK = \frac{1}{(n-1)s^3} \sum_{i=1}^n (x_i - \bar{x})^3$$

$SK < 0$ a dx rispetto
media

$SK > 0$ a sx rispetto
media

$SK = 0$ GAUSSIANA

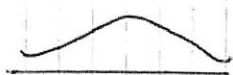
● KURTOSIS : MISURA SE LA DISTRIBUZIONE E' A PICCO O PIANA RISPETTO AD UNA NORMALE



DISTRIBUZIONE + STRETTA
PICCO + ALTO

↑ K

$K > 0$
+ ALTA



DISTRIBUZIONE + AMPIA
PICCO + BASSO
CURVA SPANCIATA

↓ K

$K < 0$
+ PIATA

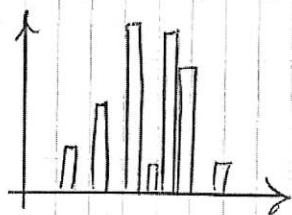
RISPETTO
A GAUSSIANA

$$K = \frac{1}{(n-1)s^4} \sum_{i=1}^n (x_i - \bar{x})^4$$

Ci sono inoltre 3 possibili rappresentazioni grafiche:

① Diagramma a barre

Gli elementi sono suddivisi in gruppi e poi ne si calcolano le frequenze

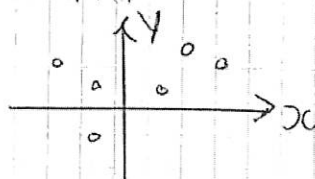


in Excel

- Data Analysis
- Istogramma
- INPUT (colonna di valori)
- Bean (intervalli 20-40-160)
- OUTPUT (legno dove)
- Legenda

② Diagramma casuale (scatter)

Ci sono 2 variabili, 1 su ascissa e 1 su ordinata e vengono inseriti i vari elementi come punti

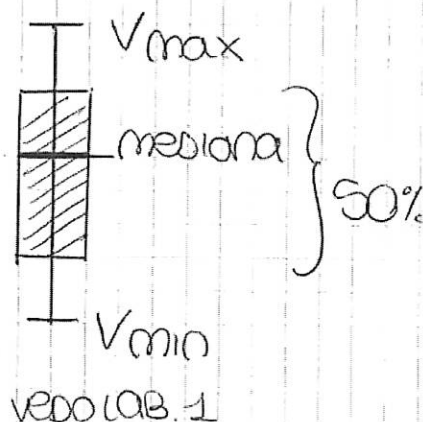


③ Diagramma a box

Il campione viene diviso in 4 percentili

È un box che contiene il 50% dei dati

Dice se la popolazione è equamente distribuita grazie alla posizione della mediana



2° PROBLEMA che relazione c'è tra i parametri della popolazione e i corrispondenti parametri del campione?

ovvero, dato un set di elementi presi da una pop, come usarlo x studiare qst pop?

- con la statistica descrittiva sappiamo come è fatto il campione
- poi si passa alla statistica inferenziale x fare delle stime e sapere come sono legati tutti i parametri visti ↓

il valore vero della popolazione si raggiunge se aumentano molto i campioni

(questo è un problema usuale quando si cerca di analizzare i dati raccolti durante un trial)

STATISTICA INFERENZIALE

SI BASA SUI CONCETTI DI:

- **PROBABILITÀ** : definita come il numero di eventi favorevoli diviso il numero totale dei possibili eventi
 $[Pr(A)]$

cioè se si verifica un evento

può essere stimata calcolando la frequenza degli eventi

- **PROBABILITÀ CONDIZIONATA** : probabilità che si verifichi A, se si è già verificato B
 $[Pr(A/B)]$

- **VARIABILI CASUALI** : variabili le cui possibili valori sono i risultati numerici di un fenomeno casuale

Sono di 2 tipi < discrete
continue

I parametri utilizzati per studiare il comportamento congiunto di 2 variabili sono:

- **COVARIANZA** = misura di come i cambiamenti in una variabile casuale sono associati con eventi di una seconda variabile casuale

cioè è il livello di variazione congiunto di 2 variabili associate linearmente
 $Cov(x, y)$

- **COEFFICIENTE DI CORRELAZIONE**: $\rho = \frac{Cov(x, y)}{\sigma_x \sigma_y}$

$$-1 < \rho < 1$$

$\rho = 0$: 2 variabili indipendenti / non correlate

$\rho = 1$: variabili $\propto$

$\rho = -1$: variabili inversamente $\propto$

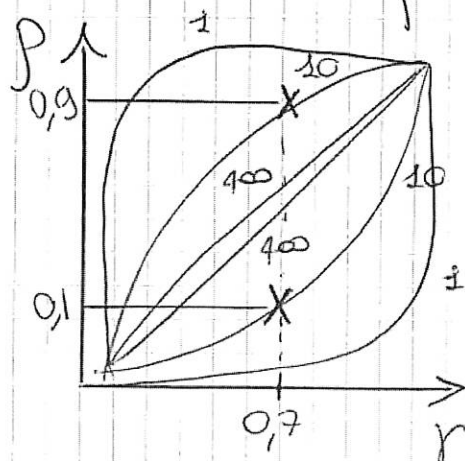
$\rho \rightarrow 0$: relazione tra le 2 variabili molto bassa

$\rho \rightarrow \pm 1$: correlazione alta! Basta usare solo una variabile su 2

⚠️ però al valore di questa stima

Soltanto, se una stima veritiera bisogna avere
 $p > 0,7$ v $p < -0,7$
 se $-0,7 < p < 0,7$, dubbi!! i punti sono molto aspri!
 Grafica pp. 13

per analizzare p si usa il **DIAGRAMMA AD OCCHIO**



ci sono 7 curve in base al numero di elementi analizzati

r è la stima della correlazione che si calcola sul campione su tot elementi

Trovato r (ex $r=0,7$) lo interseco con le curve nel numero di elementi, poi parto i 2 valori a r_x e vedo in che range sta p

$r=0,7 \rightarrow p \in [0,1; 0,9] \times 10$ elementi!
 non va bene!

$r=0,7 \rightarrow p \in [0,6-0,8] \times 100$ elementi

Intervallo ragionevole e significatività maggiore!

La numerosità del campione è quindi importante:

con maggiori elementi, la stima migliora!

"deduzioni" o "generalizzazioni"

Le inferenze vanno sempre documentate e ragionate, soprattutto quando si hanno pochi campioni! (cio non vale x la statistica descrittiva)

3° PROBLEMA Bisogna costruire dei parametri / delle variabili che siano in grado di rappresentare tutto il data set
 Come ottenere l'intero set di parametri che rappresenti il set di dati

X FORSÌ SI UTILIZZA LA **REGRESSIONE LINEARE**

= Metodo di apprendimento statistico utile ed ampiamente utilizzato

[La regressione lineare è un semplice approccio lineare utilizzato per predire una risposta quantitativa y sulla base di una singola variabile x

Sono 2 variabili:

x = V. indipendente

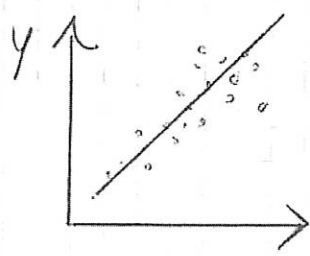
y = V. dipendente

$$y = f(x)$$

La regressione lineare assume che esista approssimativamente una relazione lineare tra x e y e che può essere rappresentata dalla retta di regressione:

$$y = \alpha + \beta x$$

α : intercetta; β : pendenza



∃ diversi tipi di fitting (Lineare, quadratico) =

sono delle equazioni che interpolano al meglio la mia dispersione di punti

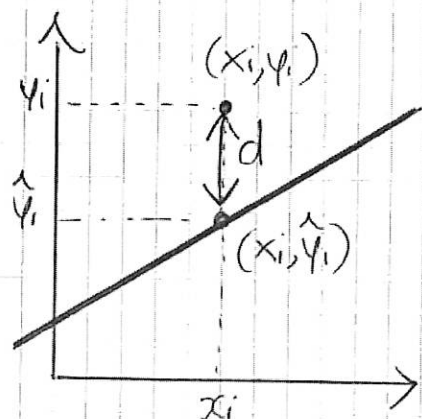
(la regressione lineare è un sottogruppo dei fitting possibili)

la varianza σ^2 (dispersione dei punti) è indice della bontà della regressione:

- I punti devono stare la maggior parte sull'interpolazione e non fuori
- altrimenti non è un modello ottimale x_i miei dati (cioè se σ^2 è troppo elevato)

x calcolare i coefficienti α e β della regressione
 Bisogna usare i dati ed il **METODO dei MINIMI QUADRATI**.
 I punti $(x_1, y_1), \dots, (x_n, y_n)$ rappresentano i dati in coppie di n -osservazioni, ciascuna delle quali consiste in una misurazione di x e y .

L'obiettivo è stimare α e β in modo che il modello lineare della retta si adatti ai dati disponibili; cioè trovare un'intercetta α ed una pendenza β in modo che la linea risultante sia la più vicina possibile ai dati.



per ogni punto ho 2 coppie di valori

(x_i, y_i) quella reale

$(x_i, \hat{y}_i)$ quella che sta sulla retta della regressione

$$\hat{y}_i = f(x_i)$$

misuriamo la distanza

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

d rappresenta l'1-esimo **Residuo**, cioè la differenza tra il valore reale osservato e il valore previsto dal modello lineare supposto.

si vuole una retta che minimizza la somma ^{dei quadrati} di tutte queste distanze $S = \sum_{i=1}^n d_i^2$

ottenendo le 2 formule

$$b = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sum_{i=1}^n (x_i - \bar{x})^2} ; a = \bar{y} - b\bar{x}$$

$(\bar{x}, \bar{y})$ è un punto che sta sempre sulla retta

Bisogna però valutare successivamente anche la **Bontà** della regressione effettuata, tramite 2 valori:

- il componente dei residui : $R_{es} = y_i - \hat{y}_i$
- il componente della regressione : $R_{eg} = \hat{y}_i - \bar{y}$

una buona regressione si ha sempre se il valore assoluto del componente della regressione è maggiore di quello dei residui (sommatoria)

$$\|\hat{y}_i - \bar{y}\| > \sum_{i=1}^n \|y_i - \hat{y}_i\|$$

(ampio) (accusa)

$$\|R_{eg}\| > \sum \|R_{es}\|$$

DIAMO 20, 4 esempi con FITTING

Buono = verde $\rightarrow \frac{1}{1} \rightarrow$ abbiamo bisogno di $d \rightarrow 0$
 Medio = Giallo
 Male = Rosso $\rightarrow \frac{1}{1} \rightarrow$ se $S = \sum d_i^2$ è grande, bisogna trovare un altro metodo

Attribuita della stima dei coeff. α e β :

la vera relazione tra x e y non è mai conosciuta x la popolazione reale.

ma la regressione lineare non toccherà mai tutti i dati reali, non si troveranno mai α e β per cui non potrà avvenire.

NOVE DEFINIZIONI:

● **DISTRIBUZIONE DI PROBABILITÀ** (funzione di probabilità)
 x una variabile discreta (casuale)

è una lista di probabilità associata con ogni valore possibile

● **Funzione della densità di probabilità**
 x variabile continue (casuale)

è una funzione tale che l'area sotto la curva di funzione di densità tra 2 punti qualsiasi a e b sia uguale alla probabilità che x cada tra a e b .
 così, l'area totale sotto la curva di funzione su tutto x è $= 1$.

TEST di IPOTESI

un'IPOTESI statistica è un'affermazione circa un parametro di una popolazione,

un TEST di IPOTESI è una procedura formale usata x accettare o rigettare un'ipotesi statistica!

ci sono 2 tipi di ipotesi statistiche che vengono contrapposte:

H_0 : ipotesi nulla
 si vuole dimostrare
 dice che le osservazioni del campione
 risultano puramente casuali