
2 variabili

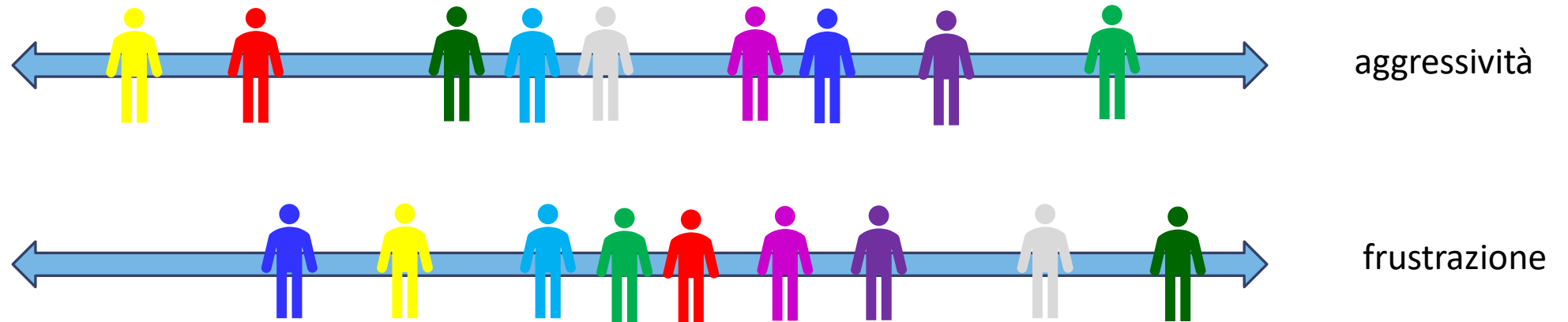
A solid blue horizontal bar spanning the entire width of the slide at the bottom.

Per imparare a leggere affermazioni apparentemente semplici da comprendere con nuove lenti,
distintive della vostra futura professione

1. La frustrazione è una determinante causale dell'aggressività.
2. Le condizioni depressive aumentano nelle persone costrette all'isolamento sociale per tre mesi.
3. L'intelligenza risente delle condizioni socio-economiche.

Between People: ci mettiamo a confronto lungo 2 variabili quantitative

L'aggressività aumenta quando aumenta la frustrazione percepita (disegno quasi-sperimentale o correlazionale).

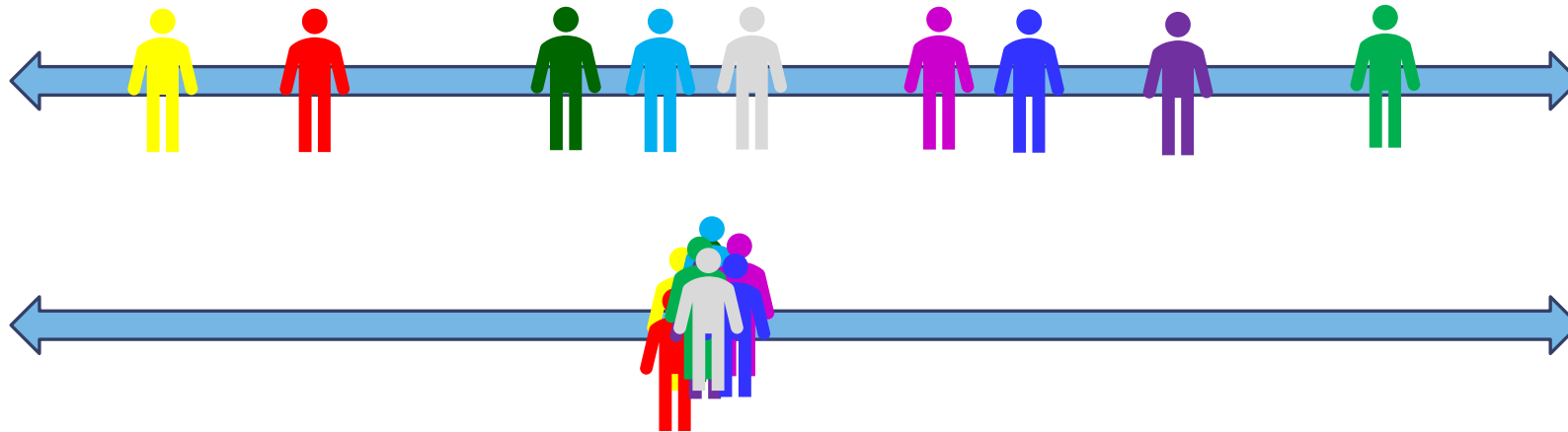


Come si collocano ciascun individuo rispetto agli altri lungo i 2 continua?

La collocazione (ordine di rango) è coerente?

Descriviamo una relazione (intensità, direzione e forma) e prevediamo l'andamento di una variabile a partire dall'altra utilizzando il coefficiente di correlazione

Between People: ci mettiamo a confronto **lungo 2 variabili quantitative**



Perché la variabilità è importante?

Prevedere una variabile quantitativa a partire da un'altra che non presenta variabilità è inefficace: persone con la stessa o quasi prestazione su una variabile hanno invece prestazioni molto diverse lungo l'altra

Between Groups (Condizioni): ci mettiamo a confronto come GRUPPI lungo 2 variabili, 1 quantitativa e 1 qualitativa



La frustrazione è una determinante causale dell'aggressività (Sperimentale).



Come si comportano mediamente le persone che appartengono all'uno vs. all'altro gruppo (es. alta vs bassa frustrazione)?

Descriviamo una relazione (intensità, direzione e forma) e prevediamo oppure spieghiamo l'andamento di una variabile a partire dall'altra (coefficiente di associazione, ANOVA btw)

NB Impareremo che le differenze nei livelli medi di Y tra due gruppi (X a due livelli) implica che Y e X sono associati o correlati

Elementi metodologici introduttivi (che riprenderemo) Per condividere un lessico comune correttamente

Relazioni dirette tra due variabili

- Co-occorrenza
- Dipendenza
- Causalità



Questi tipi di relazione non sono intercambiabili

Al di là della tecnica utilizzata (ANOVA, r di Pearson, Chi quadro, ecc),
ragionamento logico e la metodologia del disegno di ricerca ci permettono di
definire la relazione che intercorre tra 2 variabili

Elementi metodologici introduttivi

Per condividere un lessico comune correttamente

- **Co-occorrenza** (o co-variazione o correlazione): ci dice che al variare dell'una varia anche l'altra
 - è un'informazione descrittiva
 - peso/altezza; intelligenza/estroversione, auto-stima/depressione, genere/nevroticismo, orientamento politico/area cittadina
 - correlazione semplice, ANOVA btw, Chi quadro, ...

Elementi metodologici introduttivi

Per condividere un lessico comune correttamente

- **Dipendenza**: ci dice che al variare dell'una varia anche l'altra, stavolta però l'una precede logicamente l'altra variabile
 - è un'informazione descrittiva e di previsione
 - attaccamento materno/attaccamento sentimentale, intelligenza/rendimento scolastico, motivazione/successo accademico
 - correlazione semplice, ANOVA btw, Chi quadro, analisi della regressione semplice

Elementi metodologici introduttivi

Per condividere un lessico comune correttamente

- **Causalità**: ci dice che la variazione dell'una produce una variazione nell'altra
 - è un'informazione descrittiva, predittiva ed esplicativa
 - la relazione logica tra le due variabili deve essere testata con un vero esperimento
 - ANOVA btw, ... ma anche ARS o correlazione semplice sono possibili in determinati casi

Quale il ruolo di ogni variabile? Uno schema può aiutarci

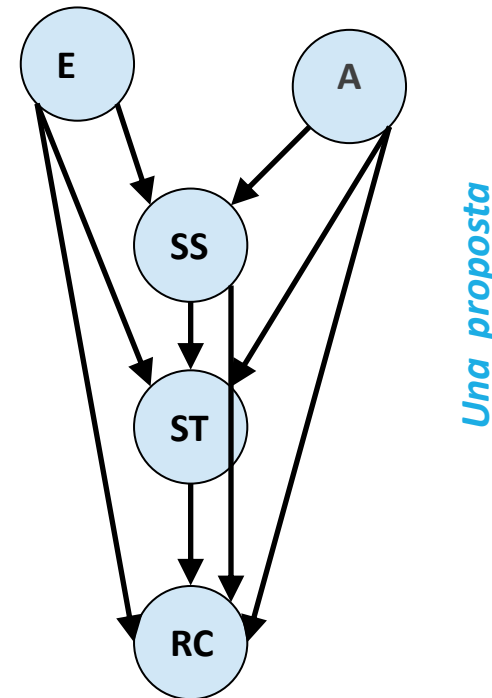
Relazioni di priorità:

Si può sostenere che una variabile precede
l'altra logicamente

o nel tempo?

Si distingue tra variabili

- esterne
- ascritte
- stati stabili
- stati transitori
- risposte comportamenti osservabili

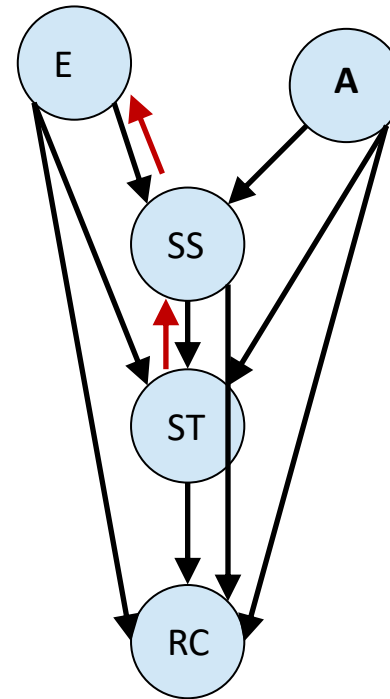


Quale il ruolo di ogni variabile? Uno schema può aiutarci

Questa tipologia di variabili permette di distinguere tra

- persona e ambiente
- stati stabili e stati temporanei
- caratteristiche intrinseche, stati e risposte comportamentali

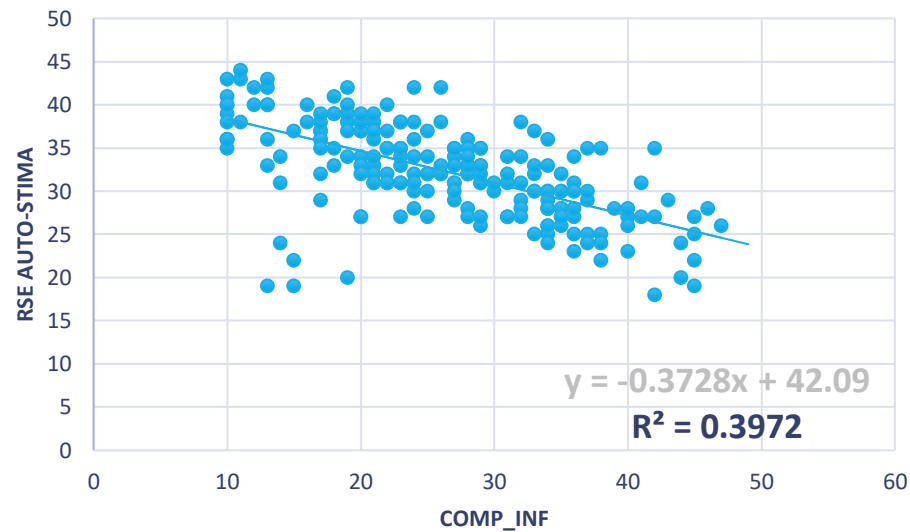
ma la direzione delle relazioni può andare anche in senso opposto, se alla base vi è un ragionamento logicamente sostenibile



Quantificare l'intensità della relazione tra due variabili continue

Useremo i nostri dati per esemplificare i concetti

Quantificare la relazione tra 2 variabili: la correlazione semplice lineare



$$codev = \sum (X_i - \bar{X})(Y_i - \bar{Y})$$

$$cov_{XY} = s_{XY} = \frac{\sum (X_i - \bar{X})(Y_i - \bar{Y})}{N-1} = codev / N - 1$$

$$r_{YX} = \frac{\sum \left(\frac{X_i - \bar{X}}{s_X} \right) \left(\frac{Y_i - \bar{Y}}{s_Y} \right)}{N-1} = \frac{cov_{XY}}{s_X s_Y}$$

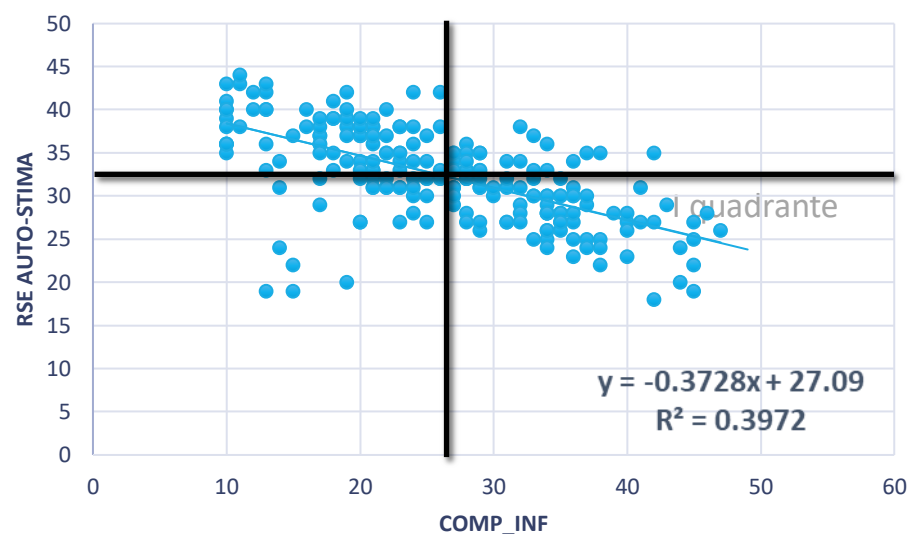
$$r_{YX} = \frac{\sum z_x z_y}{N-1}$$

*La correlazione come tecnica di analisi simmetrica,
vale a dire pone le due variabili sullo stesso piano di generalità o di priorità*

Domande via Google Form (19 ottobre)

- Potrebbe spiegare cos'è un disegno di ricerca pre-sperimentale? Quali sono i criteri per identificarlo? Questa è la definizione che ho trovato io: - non c'è campionamento da una popolazione di riferimento; - non c'è assegnazione casuale dei soggetti alle diverse condizioni; Queste sono caratteristiche sufficienti per individuarlo? Potrebbe fare alcuni esempi per avere un aiuto ad identificarli meglio?
- Un disegno sperimentale individua sempre una relazione causa-effetto? O non per forza?
- Buongiorno, è giusto affermare che nel disegno quasi-sperimentale abbiamo variabili esclusivamente di tipo quantitativo? Non mi è inoltre ben chiara la differenza tra disegno pre e quasi sperimentale. Grazie in anticipo!

Quantificare la relazione tra 2 variabili: la correlazione semplice lineare



Si traslano gli assi così che l'origine del piano coincida con i valori medi delle due variabili:

CompInf $M = 26.18 \pm 9.14$

RSE Autostima $M = 17.33 \pm 5.40$

Se applichiamo le formule, utilizzando i dati forniti in Excel, troviamo

	COMPIN	RSE		
Media	26.178	17.33		
DS	9.137	5.405		
asimmetria	0.138	-0.229		
curtosi	-0.736	-0.319		
n	214.000			
CODEV	-6628.608			
COV	-31.120			
r	-0.630		r	-0.6301
Rquadro	0.397			

Come interpretare un coefficiente di correlazione?

Significatività statistica della correlazione

Per riprendere alcuni concetti di statistica

- generalizzazione: dal campione alla popolazione
- il campione non è mai perfettamente rappresentativo della popolazione
- ipotesi statistiche: H_{nulla} e $H_{\text{alternativa}}$
- verifica mediante il calcolo di una statistica idonea sui dati del campione
- il valore della statistica rappresenta uno stimatore del parametro della popolazione
- regola di decisione: rischio che si è disposti a correre di sbagliare se si decide di rifiutare H_{nulla}

	H_0 vera	H_0 falsa
H_0 accettata	ok	errore beta
H_0 respinta	errore alfa	ok

Come interpretare un coefficiente di correlazione? Significatività statistica della correlazione

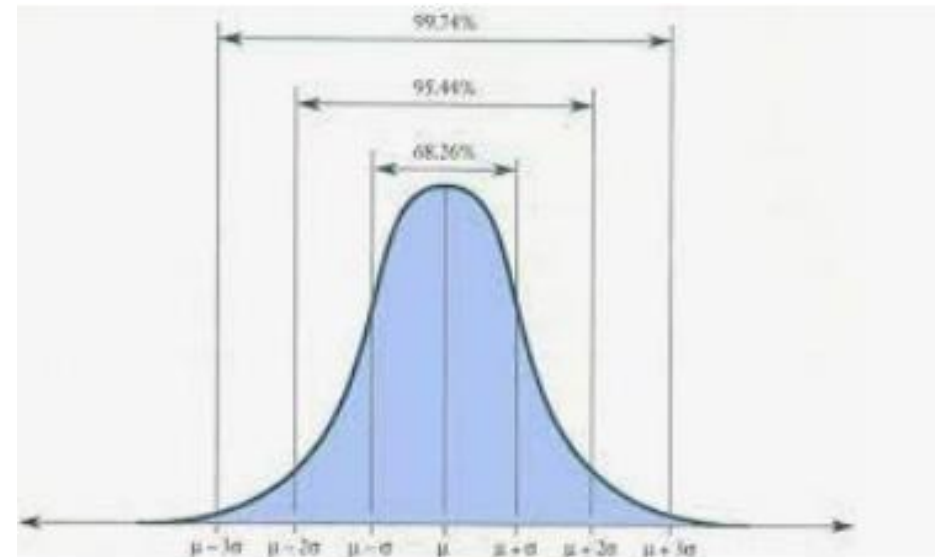
Per riprendere alcuni concetti di statistica

- il rischio è definito in termini probabilistici: definiamo la regione di rifiuto dell' H_{nulla}
- il rischio è legato all' **errore di I tipo** ed è solitamente molto piccolo (es. 5%)
- viene definito a priori da livello prescelto **alfa critico (in accordo con le proprietà dei dati e gli esiti attesi)**
- al valore della statistica osservata si associa così un grado di probabilità di osservare proprio quel valore nella popolazione **per effetto del caso**, dove però il valore atteso è definito dall' H_{nulla} (es. $\rho = 0$)
- questo grado di probabilità viene definito per mezzo delle proprietà della funzione campionaria cui appartiene la statistica (**N campione, n variabili, statistica, distribuzione campionaria della statistica**)
- distribuzione campionaria della statistica: distribuzione di frequenza con i valori della statistica calcolata sui differenti campioni della stessa numerosità, generata teoricamente prendendo infiniti campioni di dimensione n e calcolando i valori della statistica per ogni campione
- così si arriva a definire una funzione di probabilità, cioè la probabilità che quella statistica assuma proprio quel valore in quello spazio campionario

Per illustrare:
distribuzione normale

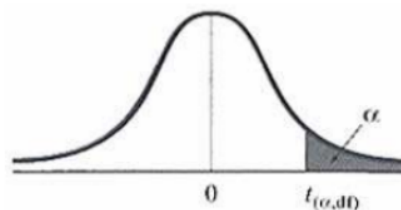
TAV. II. Aree della distribuzione normale standard tra $a = 0$ e $b > 0$

z	0	1	2	3	4	5	6	7	8	9
0.0	.0000	.0040	.0080	.0120	.0160	.0199	.0239	.0279	.0319	.0359
0.1	.0398	.0438	.0478	.0517	.0557	.0596	.0636	.0675	.0714	.0754
0.2	.0793	.0832	.0871	.0910	.0948	.0987	.1026	.1064	.1103	.1141
0.3	.1179	.1217	.1255	.1293	.1331	.1368	.1406	.1443	.1480	.1517
0.4	.1554	.1591	.1628	.1664	.1700	.1736	.1772	.1808	.1844	.1879
0.5	.1915	.1950	.1985	.2019	.2054	.2088	.2123	.2157	.2190	.2224
0.6	.2258	.2291	.2324	.2357	.2389	.2422	.2454	.2486	.2518	.2549
0.7	.2580	.2612	.2642	.2673	.2704	.2734	.2764	.2794	.2823	.2852
0.8	.2881	.2910	.2939	.2967	.2996	.3023	.3051	.3078	.3106	.3133
0.9	.3159	.3186	.3212	.3238	.3264	.3289	.3315	.3340	.3365	.3389
1.0	.3413	.3438	.3461	.3485	.3508	.3531	.3554	.3577	.3599	.3621
1.1	.3643	.3665	.3686	.3708	.3729	.3749	.3770	.3790	.3810	.3830
1.2	.3849	.3869	.3888	.3907	.3925	.3944	.3962	.3980	.3997	.4015
1.3	.4032	.4049	.4066	.4082	.4099	.4115	.4131	.4147	.4162	.4177
1.4	.4192	.4207	.4222	.4236	.4351	.4265	.4279	.4292	.4306	.4319



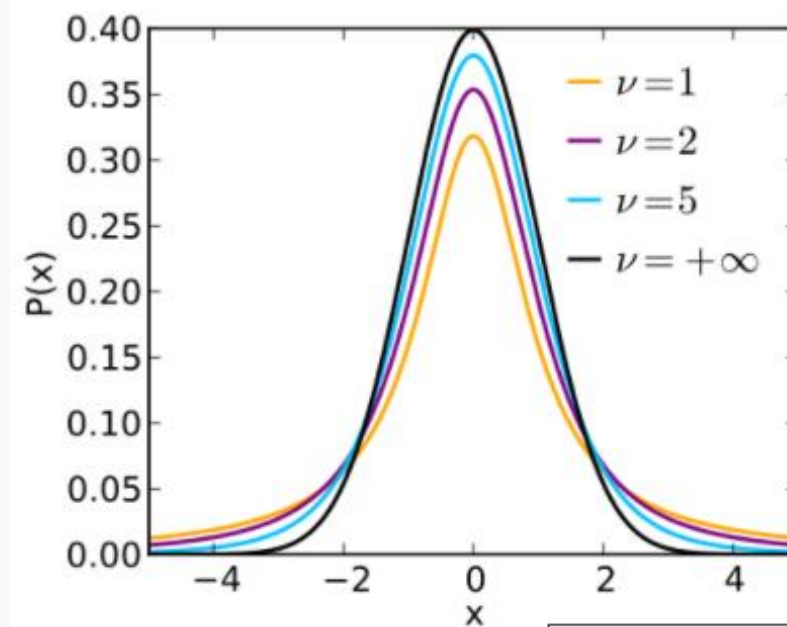
Per illustrare:

Tavola della distribuzione T di Student



Gradi di libertà	Area nella coda di destra								
	0.1	0.05	0.025	0.02	0.01	0.005	0.0025	0.001	0.0005
1	3.078	6.314	12.706	15.894	31.821	63.656	127.321	318.289	636.578
2	1.886	2.920	4.303	4.849	6.965	9.925	14.089	22.328	31.600
3	1.638	2.353	3.182	3.482	4.541	5.841	7.453	10.214	12.924
4	1.533	2.132	2.776	2.999	3.747	4.604	5.598	7.173	8.610
5	1.476	2.015	2.571	2.757	3.365	4.032	4.773	5.894	6.869
6	1.440	1.943	2.447	2.612	3.143	3.707	4.317	5.208	5.959
7	1.415	1.895	2.365	2.517	2.998	3.499	4.029	4.785	5.408
8	1.397	1.860	2.306	2.449	2.896	3.355	3.833	4.501	5.041
9	1.383	1.833	2.262	2.398	2.821	3.250	3.690	4.297	4.781
10	1.372	1.812	2.228	2.359	2.764	3.169	3.581	4.144	4.587
11	1.363	1.796	2.201	2.328	2.718	3.106	3.497	4.025	4.437
12	1.356	1.782	2.179	2.303	2.681	3.055	3.428	3.930	4.318
13	1.350	1.771	2.160	2.282	2.650	3.012	3.372	3.852	4.221
14	1.345	1.761	2.145	2.264	2.624	2.977	3.326	3.787	4.140
15	1.341	1.753	2.131	2.249	2.602	2.947	3.286	3.733	4.073
16	1.337	1.746	2.120	2.235	2.583	2.921	3.252	3.686	4.015

Funzione di densità di probabilità



la correlazione semplice lineare: verifica della significatività statistica

Per la correlazione semplice Ipotesi nulla $H_0 : \rho_{XY} = 0$

La distribuzione campionaria è quella del **t di Student**

$$t = \frac{r_{XY} \sqrt{N-2}}{\sqrt{1-r_{XY}^2}} \quad \text{con} \quad gl = n-2$$

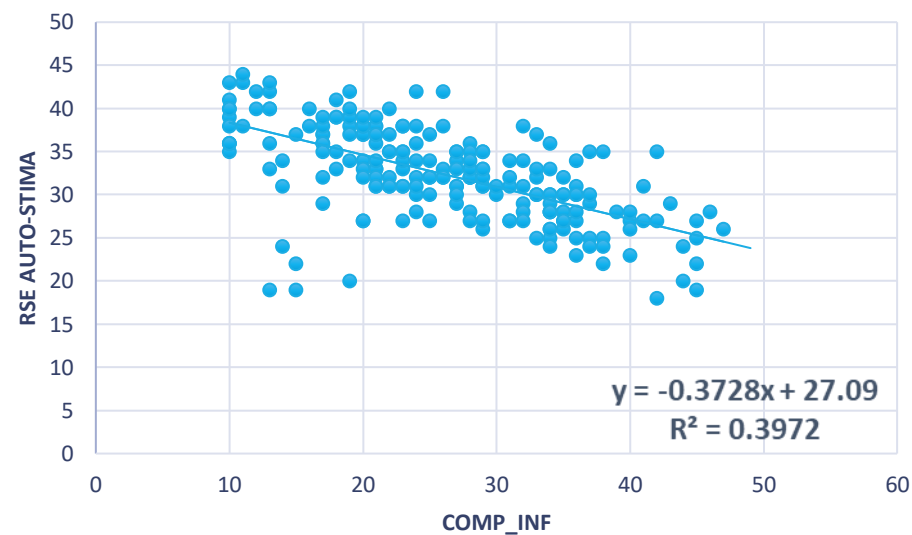
 *Quota di varianza che le 2 variabili NON condividono*

La significatività statistica è una proprietà dei dati

I gradi di libertà rappresentano il numero di possibilità che i dati che compongono un campione hanno di variare liberamente.
Nella distribuzione dei valori di X e di Y osservati, un valore di X e un valore di Y sono vincolati alle rispettive medie

variamo i valori nella formula di t	
aumentiamo N	-14.0592986
abbassiamo r	-8.39643197

la correlazione semplice lineare: verifica della significatività statistica



Se applichiamo le formule, utilizzando i dati forniti in Excel, troviamo che

$t = -11.82$

$gl = 214-1$

$p < .001$

NB. Se variamo N oppure r, il valore di t cambia

aumentiamo N (= 300)	-14.01
abbassiamo r = -.20)	-9.36

Oltre alla significatività statistica, è essenziale la significatività sostanziale data dalla grandezza dell’effetto rilevato

Un esempio da Jamovi

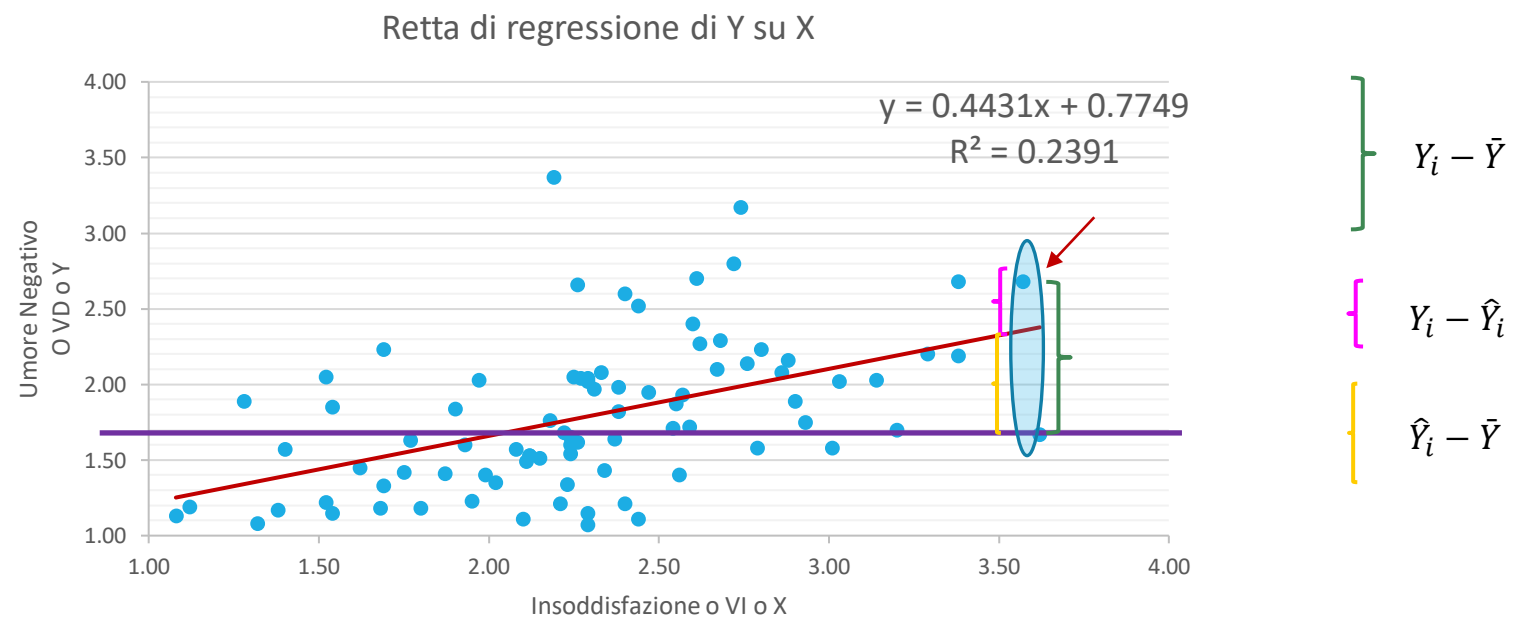
Correlation Matrix

		COMPIN	RSE
COMPIN	Pearson's r	—	
	p-value	—	
	95% CI Upper	—	
	95% CI Lower	—	
	N	—	
RSE	Pearson's r	-0.630	—
	p-value	< .001	—
	95% CI Upper	-0.542	—
	95% CI Lower	-0.705	—
	N	214	—

Regressione semplice lineare

- quantifica l'intensità della relazione lineare tra due variabili proprio come una correlazione semplice lineare
- permette di rappresentare graficamente una correlazione lineare (o anche non lineare) tra due variabili
- rende una relazione logicamente simmetrica in una relazione tecnicamente asimmetrica, obbligando a indicare una variabile come VD o Y e l'altra come VI o X
- permette di prevedere, con un margine di errore, l'andamento di una variabile a partire dall'altra

Regressione semplice lineare: il metodo dei minimi quadrati



Regressione semplice lineare: dalla rappresentazione grafica all'equazione di regressione

metodo dei minimi quadrati $\sum (Y_i - \hat{Y}_i)^2 = \min$

nella popolazione, l'eq che descrive la relazione tra X e Y è la seguente:

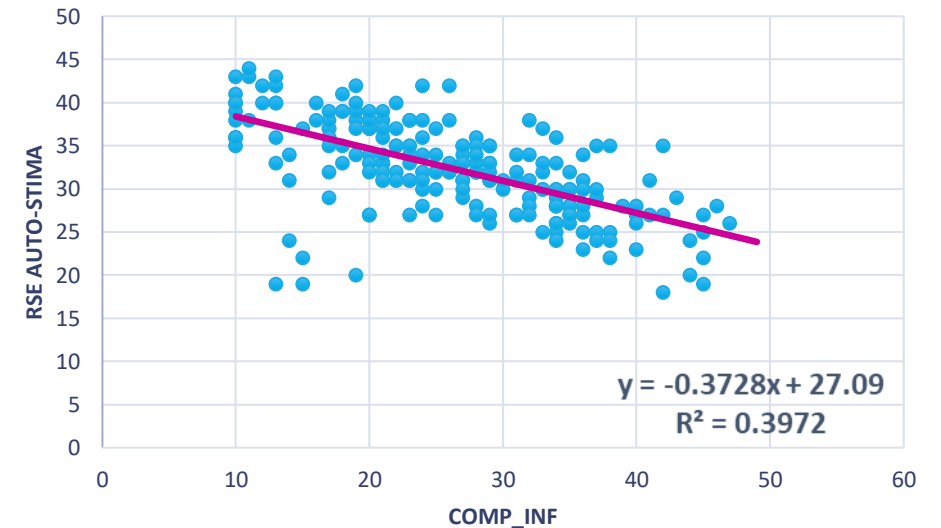
$$Y = \alpha + \beta X + \varepsilon$$

nel campione, l'eq di previsione è la seguente

$$\hat{Y}_i = a + bX_i$$

$$Y_i = a + bX_i + e$$

$\hat{Y}$ è una trasformazione lineare di X



Regressione semplice lineare:
dalla rappresentazione grafica all'equazione di regressione

metodo dei minimi quadrati $\sum (Y_i - \hat{Y}_i)^2 = \min$

nella popolazione, l'eq che descrive la relazione tra X e Y è la seguente:

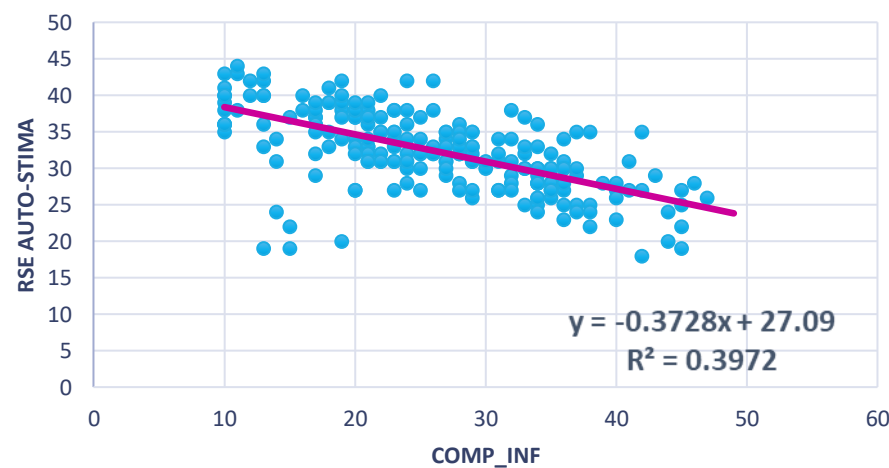
$$Y = \alpha + \beta X + \varepsilon$$

nel campione, l'eq di previsione è la seguente

$$\hat{Y}_i = a + bX_i$$

$$Y_i = a + bX_i + e$$

$\hat{Y}$ è una trasformazione lineare di X



Regressione semplice lineare

$$\hat{Y}_i = a_{YX} + b_{YX} X_i$$

a_{YX} rappresenta una stima campionaria di α_{YX}

- il valore di $\hat{Y}$ quando $X = 0$ (e l'asse X incrocia l'asse Y)
- ed è dato dalla seguente formula, per la stima dei minimi quadrati:
$$a_{YX} = \bar{Y} - b_{YX} \bar{X}$$
- per lo “strabismo” della tecnica di regressione $a_{YX} \neq a_{XY}$
- se la VI viene centrata intorno alla media (**centering**: $X_i - \bar{X}$ trasformare i punteggi osservati per la VI in punteggi espressi in deviazione dalla media) l'interpretazione dei coefficienti a e b è più semplice e chiara rispetto ai dati osservati, quando non sono già standardizzati

Regressione semplice lineare

$$\hat{Y}_i = a_{YX} + b_{YX} X_i$$

b_{YX} rappresenta

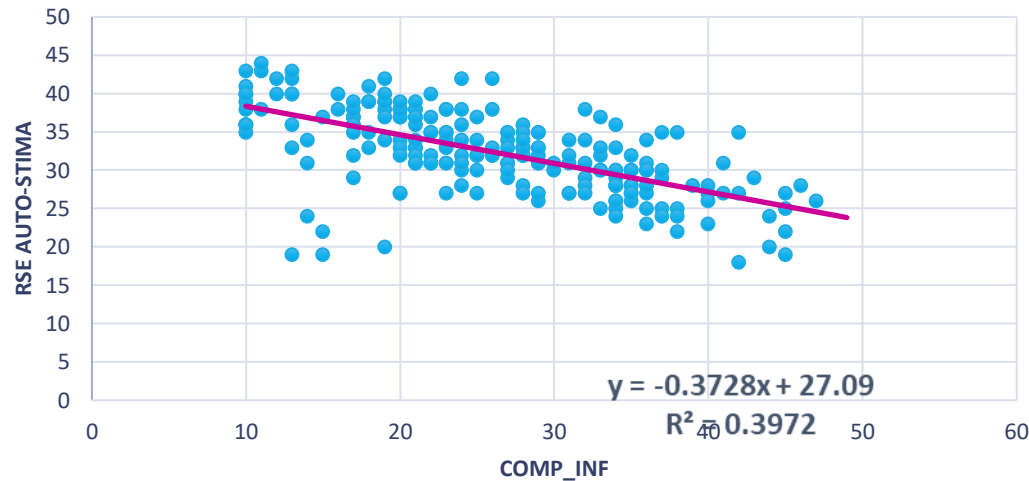
- una stima campionaria di β_{YX}
- la variazione in $\hat{Y}$ al variare di 1 unità in X
- una misura di associazione tra Y e X, non standardizzata, la cui formula per la stima dei minimi quadrati è

$$b_{YX} = \frac{\sum(X_i - \bar{X})(Y_i - \bar{Y})}{\sum(X_i - \bar{X})^2} = \frac{\frac{\sum(X_i - \bar{X})(Y_i - \bar{Y})}{N-1}}{\frac{\sum(X_i - \bar{X})^2}{N-1}} = \frac{s_{XY}}{s_X^2}$$

$$b_{YX} = \frac{s_{XY}}{s_X^2} \times \frac{s_Y}{s_Y} = \frac{s_{XY}}{s_X s_Y} \times \frac{s_Y}{s_X} = r_{XY} \times \frac{s_Y}{s_X}$$

$$\frac{s_{XY}}{s_X s_Y} = r_{xy}$$

Regressione semplice lineare



Se ritorniamo ai nostri dati in excel,
l'eq di regressione secondo il metodo
dei minimi quadrati è la seguente:

$$Y' = 27.09 + (-0.373 * X_i)$$

A partire da questa eq. è possibile stimare
Y (RSE) a partire CompInf,
i valori attesi si collocano lungo la retta
di regressione (in fuxia),
con un margine di **errore di stima** dato da $Y - Y'$

Questo scarto rappresenta quella parte del punteggio
di Y osservato che non si associa e non può essere
prevista a partire da X, un residuo che più in là
riponderemo ...

Regressione semplice lineare

- se X e Y sono standardizzate, allora il coefficiente std è rappresentato da $\hat{\beta}$ che nella regressione lineare semplice coincide con il valore della correlazione semplice

$$\hat{\beta}_{YX} = r_{YX} = r_{XY}$$

$$r_{XY} \times \frac{s_Y}{s_X}$$

$$\hat{Y}_i = a_{YX} + b_{YX} X_i$$

- per lo “strabismo” della tecnica di regressione $b_{YX} \neq b_{XY}$

Regressione semplice lineare

- a_{YX} e b_{YX} seguono una distribuzione campionaria t con valore atteso
- per verificare la significatività di a_{YX} e b_{YX} osservati, si testano le ipotesi nulle

$$H_0 : \alpha_{YX} = 0 \quad t = (a_{YX} - H_0) / SE_{a_{YX}} \quad \text{con } gl = n - 2$$

$$H_0 : \beta_{YX} = 0 \quad t = (b_{YX} - H_0) / SE_{b_{YX}} \quad \text{con } gl = n - 2$$

(già CI sono informativi rispetto a diverse possibili ipotesi nulle)

Linear Regression

Jamovi

Model Fit Measures

Model	R	R ²	Adjusted R ²	Overall Model Test			
				F	df1	df2	p
1	0.630	0.397	0.394	140	1	212	< .001

Omnibus ANOVA Test

	Sum of Squares	df	Mean Square	F	p
COMPIN	2471	1	2471.1	140	< .001
Residuals	3750	212	17.7		

Note. Type 3 sum of squares

Model Coefficients - RSE

Predictor	Estimate	SE	95% Confidence Interval		t	p	Stand. Estimate
			Lower	Upper			
Intercept	27.090	0.8743	25.367	28.814	31.0	< .001	
COMPIN	-0.373	0.0315	-0.435	-0.311	-11.8	< .001	-0.630

Regressione semplice lineare

errore standard del coefficiente b_{YX}

$$SE_{b_{YX}} = \frac{s_y}{s_x} \sqrt{\frac{1 - r_{YX}^2}{n - 2}} \quad gl (n - k - 1)$$

e intervallo di confidenza $SE_{b_{YX}} t_{crit} \pm b_{YX}$

Sempre per i nostri dati

SE b = 0.032

95% IC intorno a b = -0.373

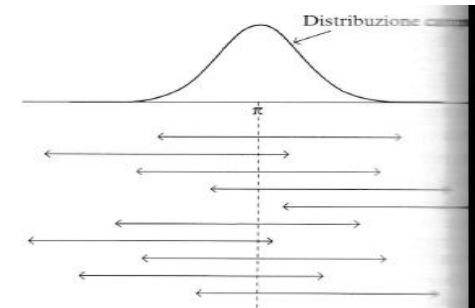
sono -0.435 e -0.311

Nell'intervallo NON è compreso lo 0

breve ripasso: l'errore standard di uno stimatore

Errore standard di uno stimatore: corrisponde alla SD della distribuzione campionaria dello stimatore, è pertanto un indicatore di variabilità o imprecisione dello stimatore campionario rispetto al valore della popolazione; tanto minore quanto maggiore è N .

Intervallo di Confidenza: «confidenza 95%» significa che se ripetessimo la stessa indagine per 100 volte con gli stessi metodi (ma su 100 campioni diversi), probabilmente otterremmo ogni volta una stima diversa e intorno a questa diversi IC; **tuttavia, il vero valore della popolazione sarebbe all'interno di intervalli di confidenza 95 volte su 100**



Il coefficiente di determinazione

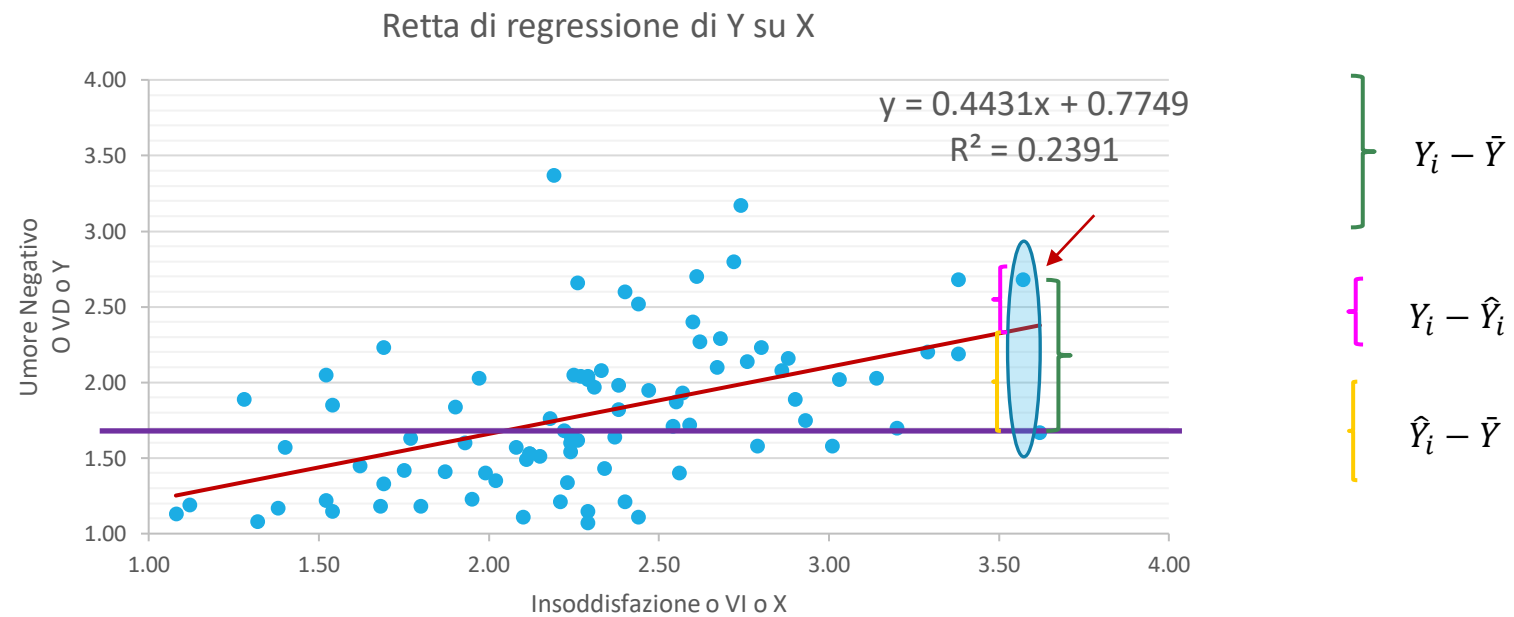
Il coefficiente di correlazione elevato al quadrato R^2 rappresenta un indice quantitativo di RPE o riduzione proporzionale dell'errore di stima

$$RPE = \frac{e_{\text{SENZA REGOLA DI DECISIONE}} - e_{\text{CON REGOLA DI DECISIONE}}}{e_{\text{SENZA REGOLA DI DECISIONE}}}$$

$$r_{YX}^2 = \frac{\sum (Y_i - \bar{Y})^2 - \sum (Y_i - \hat{Y}_i)^2}{\sum (Y_i - \bar{Y})^2} = \frac{\sum (\hat{Y}_i - \bar{Y})^2}{\sum (Y_i - \bar{Y})^2} = r_{Y\hat{Y}}^2$$

Eta quadro = DEV tra / DEV tot

$$r_{YX}^2 = \frac{\sum(Y_i - \bar{Y})^2 - \sum(Y_i - \hat{Y}_i)^2}{\sum(Y_i - \bar{Y})^2} = \frac{\sum(\hat{Y}_i - \bar{Y})^2}{\sum(Y_i - \bar{Y})^2} = r_{Y\hat{Y}}^2$$



Il coefficiente di determinazione

$$\begin{aligned} RPE &= \frac{\sum(Y - \bar{Y})^2 - \sum(Y - \hat{Y})^2}{\sum(Y - \bar{Y})^2} = \frac{\sum(\hat{Y} - \bar{Y})^2}{\sum(Y - \bar{Y})^2} = \\ &= \frac{\sum(\hat{Y} - \bar{Y})^2 / N - 1}{\sum(Y - \bar{Y})^2 / N - 1} = \frac{s_{\hat{Y}}^2}{s_Y^2} = r_{XY}^2 \end{aligned}$$

Se applichiamo queste formule ai nostri dati in Excel

RPE come $R^2 = 4,95/20,71 = 0,239$

RPE come $R^2 = (4,95/85)/(20,71/85) = 0,239$

r_{XY}^2

come rapporto tra varianza spiegata dalla regressione (dagli stimatori) e varianza totale di Y

$$F = \frac{\text{dev. regress}/k}{\text{dev. residua}/(N - k - 1)}$$

Il coefficiente di determinazione

La verifica della significatività statistica R^2

Ipotesi nulla $R^2 = 0$

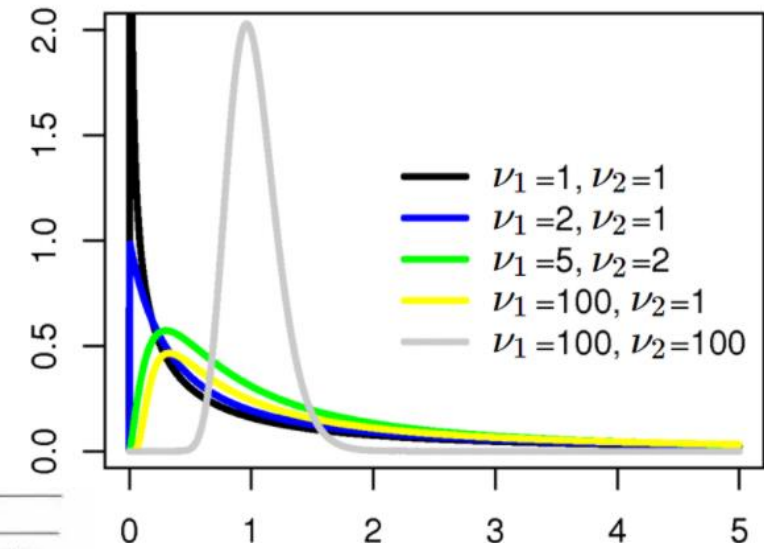
Distribuzione campionaria: F di Fisher

$$F = \frac{\text{dev. regress}/k}{\text{dev. residua}/(N - k - 1)}$$

k = numero di stimatori o VI

La decisione viene presa in base alla regola per cui $p \leq \alpha$.

Per illustrare: F di Fisher



		Gradi di libertà al numeratore																						
		1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	20	24	30	40	60	100	200	+ ∞
Gradi di libertà al denominatore	1	161,45	199,50	215,71	224,58	230,16	233,99	236,77	238,88	240,54	241,88	242,98	243,90	244,69	245,36	245,95	248,02	249,05	250,10	251,14	252,20	253,04	253,68	254,31
	2	18,51	19,00	19,16	19,25	19,30	19,33	19,35	19,37	19,38	19,40	19,40	19,41	19,42	19,42	19,43	19,45	19,45	19,46	19,47	19,48	19,49	19,49	19,50
	3	10,13	9,55	9,28	9,12	9,01	8,94	8,89	8,85	8,81	8,79	8,76	8,74	8,73	8,71	8,70	8,66	8,64	8,62	8,59	8,57	8,55	8,54	8,53
	4	7,71	6,94	6,59	6,39	6,26	6,16	6,09	6,04	6,00	5,96	5,94	5,91	5,89	5,87	5,86	5,80	5,77	5,75	5,72	5,69	5,66	5,65	5,63
	5	6,61	5,79	5,41	5,19	5,05	4,95	4,88	4,82	4,77	4,74	4,70	4,68	4,66	4,64	4,62	4,56	4,53	4,50	4,46	4,43	4,41	4,39	4,37
	6	5,99	5,14	4,76	4,53	4,39	4,28	4,21	4,15	4,10	4,06	4,03	4,00	3,98	3,96	3,94	3,87	3,84	3,81	3,77	3,74	3,71	3,69	3,67
	7	5,59	4,74	4,35	4,12	3,97	3,87	3,79	3,73	3,68	3,64	3,60	3,57	3,55	3,53	3,51	3,44	3,41	3,38	3,34	3,30	3,27	3,25	3,23
	8	5,32	4,46	4,07	3,84	3,69	3,58	3,50	3,44	3,39	3,35	3,31	3,28	3,26	3,24	3,22	3,15	3,12	3,08	3,04	3,01	2,97	2,95	2,93
	9	5,12	4,26	3,86	3,63	3,48	3,37	3,29	3,23	3,18	3,14	3,10	3,07	3,05	3,03	3,01	2,94	2,90	2,86	2,83	2,79	2,76	2,73	2,71
	10	4,96	4,10	3,71	3,48	3,33	3,22	3,14	3,07	3,02	2,98	2,94	2,91	2,89	2,86	2,85	2,77	2,74	2,70	2,66	2,62	2,59	2,56	2,54
	11	4,84	3,98	3,59	3,36	3,20	3,09	3,01	2,95	2,90	2,85	2,82	2,79	2,76	2,74	2,72	2,65	2,61	2,57	2,53	2,49	2,46	2,43	2,40
	12	4,75	3,89	3,49	3,26	3,11	3,00	2,91	2,85	2,80	2,75	2,72	2,69	2,66	2,64	2,62	2,54	2,51	2,47	2,43	2,38	2,35	2,32	2,30
	13	4,67	3,81	3,41	3,18	3,03	2,92	2,83	2,77	2,71	2,67	2,63	2,60	2,58	2,55	2,53	2,46	2,42	2,38	2,34	2,30	2,26	2,23	2,21
	14	4,60	3,74	3,34	3,11	2,96	2,85	2,76	2,70	2,65	2,60	2,57	2,53	2,51	2,48	2,46	2,39	2,35	2,31	2,27	2,22	2,19	2,16	2,13
	15	4,54	3,68	3,29	3,06	2,90	2,79	2,71	2,64	2,59	2,54	2,51	2,48	2,45	2,42	2,40	2,33	2,29	2,25	2,20	2,16	2,12	2,10	2,07
20	4,35	3,49	3,10	2,87	2,71	2,60	2,51	2,45	2,39	2,35	2,31	2,28	2,25	2,22	2,20	2,12	2,08	2,04	1,99	1,95	1,91	1,88	1,84	
24	4,26	3,40	3,01	2,78	2,62	2,51	2,42	2,36	2,30	2,25	2,22	2,18	2,15	2,13	2,11	2,03	1,98	1,94	1,89	1,84	1,80	1,77	1,73	
30	4,17	3,32	2,92	2,69	2,53	2,42	2,33	2,27	2,21	2,16	2,13	2,09	2,06	2,04	2,01	1,93	1,89	1,84	1,79	1,74	1,70	1,66	1,62	
40	4,08	3,23	2,84	2,61	2,45	2,34	2,25	2,18	2,12	2,08	2,04	2,00	1,97	1,95	1,92	1,84	1,79	1,74	1,69	1,64	1,59	1,55	1,51	
60	4,00	3,15	2,76	2,53	2,37	2,25	2,17	2,10	2,04	1,99	1,95	1,92	1,89	1,86	1,84	1,75	1,70	1,65	1,59	1,53	1,48	1,44	1,39	
100	3,94	3,09	2,70	2,46	2,31	2,19	2,10	2,03	1,97	1,93	1,89	1,85	1,82	1,79	1,77	1,68	1,63	1,57	1,52	1,45	1,39	1,34	1,28	
200	3,89	3,04	2,65	2,42	2,26	2,14	2,06	1,98	1,93	1,88	1,84	1,80	1,77	1,74	1,72	1,62	1,57	1,52	1,46	1,39	1,32	1,26	1,19	
+ ∞	3,84	3,00	2,61	2,37	2,21	2,10	2,01	1,94	1,88	1,83	1,79	1,75	1,72	1,69	1,67	1,57	1,52	1,46	1,39	1,32	1,24	1,17	1,01	

Linear Regression

Jamovi

Model Fit Measures

Model	R	R ²	Adjusted R ²	Overall Model Test			
				F	df1	df2	p
1	0.630	0.397	0.394	140	1	212	< .001

Omnibus ANOVA Test

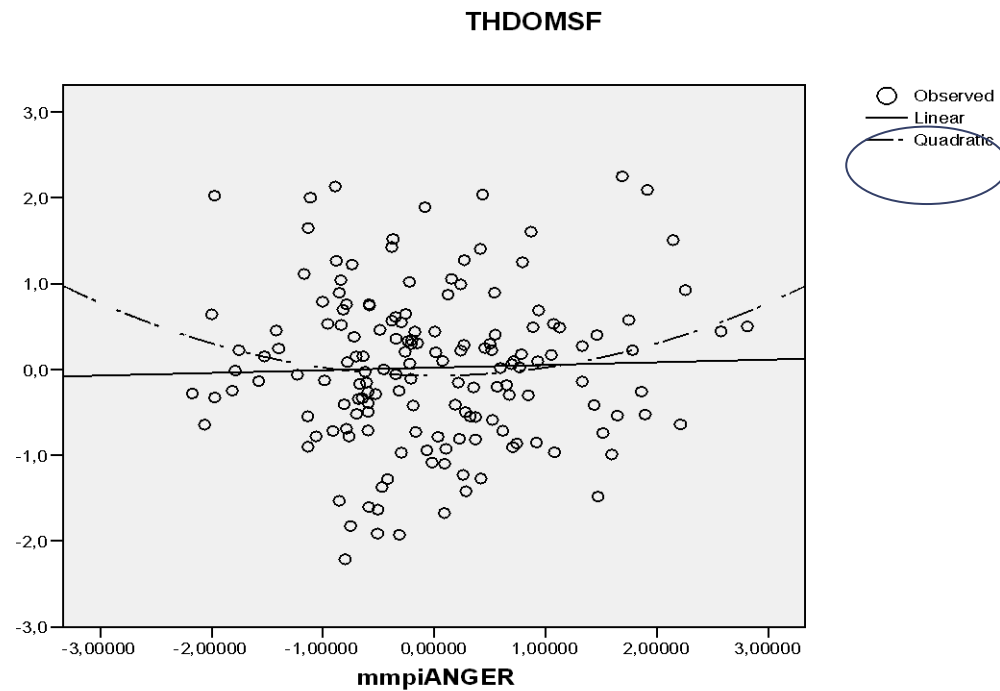
	Sum of Squares	df	Mean Square	F	p
COMPIN	2471	1	2471.1	140	< .001
Residuals	3750	212	17.7		

Note. Type 3 sum of squares

Model Coefficients - RSE

Predictor	Estimate	SE	95% Confidence Interval		t	p	Stand. Estimate
			Lower	Upper			
Intercept	27.090	0.8743	25.367	28.814	31.0	< .001	
COMPIN	-0.373	0.0315	-0.435	-0.311	-11.8	< .001	-0.630

Correlazioni non lineari



In questo esempio, la relazione è non-lineare
pertanto una correlazione semplice lineare come
quella
considerata finora indica $r = 0$
se invece si calcola una correlazione non-lineare
ma quadratica allora $r > 0$

Ancora sulla RSL: il caso di una variabile indipendente dicotomica

L'ANOVA è un caso particolare del modello lineare generale ...

ANOVA e RSL ci danno stesse informazioni

confrontiamo i risultati nel caso di un VI dicotomica codificata come «dummy», vale a dire assegnando 0 a un gruppo, detto di riferimento, e 1 all'altro

Per i nostri dati in Jamovi, VD = Humor aggressivo e VI = Genere (dicotomico)

Linear Regression

Model Fit Measures

Model	R	R ²
1	0.194	0.0376

Ancora sulla RSL: il caso di una variabile indipendente dicotomica

Omnibus ANOVA Test

	Sum of Squares	df	Mean Square	F	p
Qualeilsuogenere	356	1	355.9	7.04	0.009
Residuals	9099	180	50.6		

Note. Type 3 sum of squares

Model Coefficients - HUMOR_AGGRESSIVE

Predictor	Estimate	SE	t	p	Stand. Estimate
Intercept	26.40	1.10	24.07	< .001	
Qualeilsuogenere	-3.32	1.25	-2.65	0.009	-0.194

2 variabili: VD quantitative e VI categoriale

ANOVA btw

Incominciamo con un disegno sperimentale classico

Obiettivo: verificare l'ipotesi secondo cui il trattamento (VI) influenza causalmente le reazioni dei Ss alla prova e che costituiscono il post-test (VD)

	Pre-test	Trattamento	Post-test	
GrSperim (R)	si	si	si	→
GrContro (R)	si	no	si	→

*Medie
a confronto*

Si possono controllare:

- effetto della selezione, della storia attuale, della maturazione, ... grazie all'equivalenza tra i gruppi, dovuta all'assegnazione randomizzata dei Ss
- effetto del pre-test sul post-test, poiché tutti i gruppi sono stati sottoposti a pre-test

2 variabili: VD quantitative e VI categoriale ANOVA btw

Incominciamo con un disegno sperimentale classico

Obiettivo: verificare l'ipotesi secondo cui il trattamento (VI) influenza causalmente le reazioni dei Ss alla prova e che costituiscono il post-test (VD)

	Trattamento	Post-test	
GrSperim (R)	si	si	→ Medie a confronto
GrContro (R)	no	si	

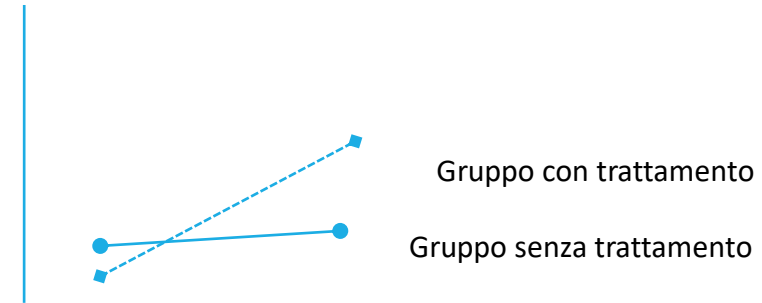
Grazie alla randomizzazione, presenta gli stessi vantaggi (e limiti) del disegno classico, ma non è possibile verificare direttamente equivalenza dei gruppi, effetti della maturazione, ... Disegno ampiamente usato.

2 variabili: VD quantitative e VI categoriale ANOVA btw

- è possibile verificare l'equivalenza statistica tra i gruppi non sperimentali
- tuttavia anche quando verificata non si può escludere del tutto che non la VI ma altre variabili *confounded* abbiano agito sugli esiti del post-test
- pertanto si verifica l'efficacia della VI trattamento in termini differenze pre-post test osservate tra i gruppi studiati

	Pre-test	Trattamento	Post-test	Differenza
Gruppo A (Nr)	si	si	si	pre-post
Gruppo B (Nr)	si	no	si	pre-post

- il *cross-over effect* garantisce più di altri risultati l'efficacia del trattamento



Logica dell'**ANOVA**: scomposizione della variabilità delle reazioni dei Ss (devianza) in un disegno classico a 1 fattore btw (o VI)

1. è possibile calcolare una **media generale M_{GEN}** della VD (effetto generale), per l'intero campione (GrSp + GrCo): ogni soggetto tuttavia può discostarsi da questa prestazione media e tale scostamento è detto scarto

la somma dei quadrati di questi scarti rappresenta la **DEVIANZA TOTALE**, un **indice quantitativo della variabilità attorno alla media**

Logica dell'ANOVA: scomposizione della variabilità delle reazioni dei Ss (devianza) in un disegno classico a 1 fattore btw (o VI)

2. è possibile calcolare una media della VD anche per ogni gruppo (effetto della VI o del trattamento, M_{gruppo_j}): se assumo che ogni soggetto appartenente ad un gruppo ha un punteggio pari alla media del suo gruppo e calcolo la differenza tra punteggio medio generale e punteggio medio di gruppo, ottengo uno scarto

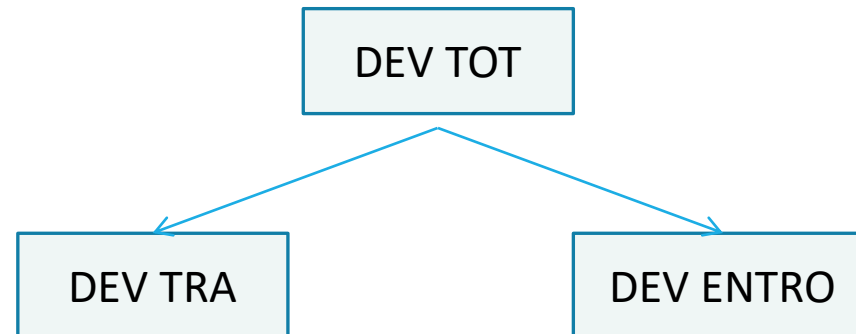
la somma dei quadrati di questi scarti delle medie di gruppo rispetto alla media generale rappresenta la DEVIANZA TRA i gruppi, un indice quantitativo della variabilità tra i gruppi, come variano al di là dell'andamento generale

Logica dell'ANOVA: scomposizione della variabilità delle reazioni dei Ss (devianza) in un disegno classico a 1 fattore btw (o VI)

3. è possibile verificare come ogni soggetto si discosti dalla media del suo gruppo (**effetto residuo**, non attribuibile al trattamento, alla VI):

la somma dei quadrati di questi scarti rappresenta la **DEVIANZA ENTRO i gruppi**, un indice quantitativo della variabilità all'interno di ogni gruppo; è una **variabilità residuale**, non spiegabile con la VI (errore casuale)

Si può dimostrare che $DEV\ TOT = DEV\ TRA + DEV\ ENTRO$ ossia



Logica dell'ANOVA: scomposizione della variabilità delle reazioni dei Ss (devianza) in un disegno classico a 1 fattore btw (o VI)

$$y_{ij} = \mu + \alpha + \varepsilon$$

$$y_{ij} = M_{GEN} + (M_{gruppo_j} - M_{GEN}) + (y_{ij} - M_{gruppo_j})$$

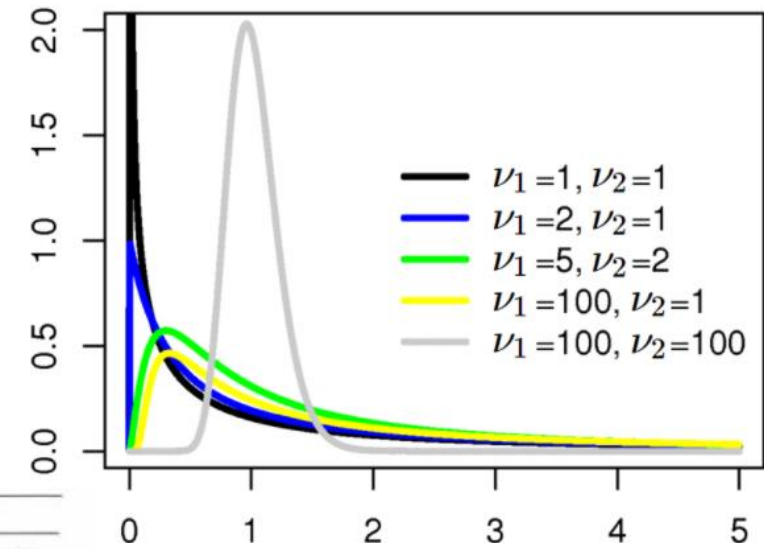
In termini di scostamento dalla media generale:

$$(y_{ij} - M_{GEN}) = (M_{gruppo_j} - M_{GEN}) + (y_{ij} - M_{gruppo_j})$$

ovvero in termini di scomposizione della devianza totale

$$\sum \sum (y_{ij} - M_{GEN})^2 = \sum \sum (M_{gruppo_j} - M_{GEN})^2 + \sum \sum (y_{ij} - M_{gruppo_j})^2$$

Per illustrare: F di Fisher



Gradi di libertà al numeratore																								
	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	20	24	30	40	60	100	200	+ ∞	
Gradi di libertà al denominatore	1	161,45	199,50	215,71	224,58	230,16	233,99	236,77	238,88	240,54	241,88	242,98	243,90	244,69	245,36	245,95	248,02	249,05	250,10	251,14	252,20	253,04	253,68	254,31
	2	18,51	19,00	19,16	19,25	19,30	19,33	19,35	19,37	19,38	19,40	19,40	19,41	19,42	19,42	19,43	19,45	19,45	19,46	19,47	19,48	19,49	19,49	19,50
	3	10,13	9,55	9,28	9,12	9,01	8,94	8,89	8,85	8,81	8,79	8,76	8,74	8,73	8,71	8,70	8,66	8,64	8,62	8,59	8,57	8,55	8,54	8,53
	4	7,71	6,94	6,59	6,39	6,26	6,16	6,09	6,04	6,00	5,96	5,94	5,91	5,89	5,87	5,86	5,80	5,77	5,75	5,72	5,69	5,66	5,65	5,63
	5	6,61	5,79	5,41	5,19	5,05	4,95	4,88	4,82	4,77	4,74	4,70	4,68	4,66	4,64	4,62	4,56	4,53	4,50	4,46	4,43	4,41	4,39	4,37
	6	5,99	5,14	4,76	4,53	4,39	4,28	4,21	4,15	4,10	4,06	4,03	4,00	3,98	3,96	3,94	3,87	3,84	3,81	3,77	3,74	3,71	3,69	3,67
	7	5,59	4,74	4,35	4,12	3,97	3,87	3,79	3,73	3,68	3,64	3,60	3,57	3,55	3,53	3,51	3,44	3,41	3,38	3,34	3,30	3,27	3,25	3,23
	8	5,32	4,46	4,07	3,84	3,69	3,58	3,50	3,44	3,39	3,35	3,31	3,28	3,26	3,24	3,22	3,15	3,12	3,08	3,04	3,01	2,97	2,95	2,93
	9	5,12	4,26	3,86	3,63	3,48	3,37	3,29	3,23	3,18	3,14	3,10	3,07	3,05	3,03	3,01	2,94	2,90	2,86	2,83	2,79	2,76	2,73	2,71
	10	4,96	4,10	3,71	3,48	3,33	3,22	3,14	3,07	3,02	2,98	2,94	2,91	2,89	2,86	2,85	2,77	2,74	2,70	2,66	2,62	2,59	2,56	2,54
	11	4,84	3,98	3,59	3,36	3,20	3,09	3,01	2,95	2,90	2,85	2,82	2,79	2,76	2,74	2,72	2,65	2,61	2,57	2,53	2,49	2,46	2,43	2,40
	12	4,75	3,89	3,49	3,26	3,11	3,00	2,91	2,85	2,80	2,75	2,72	2,69	2,66	2,64	2,62	2,54	2,51	2,47	2,43	2,38	2,35	2,32	2,30
	13	4,67	3,81	3,41	3,18	3,03	2,92	2,83	2,77	2,71	2,67	2,63	2,60	2,58	2,55	2,53	2,46	2,42	2,38	2,34	2,30	2,26	2,23	2,21
	14	4,60	3,74	3,34	3,11	2,96	2,85	2,76	2,70	2,65	2,60	2,57	2,53	2,51	2,48	2,46	2,39	2,35	2,31	2,27	2,22	2,19	2,16	2,13
	15	4,54	3,68	3,29	3,06	2,90	2,79	2,71	2,64	2,59	2,54	2,51	2,48	2,45	2,42	2,40	2,33	2,29	2,25	2,20	2,16	2,12	2,10	2,07
20	4,35	3,49	3,10	2,87	2,71	2,60	2,51	2,45	2,39	2,35	2,31	2,28	2,25	2,22	2,20	2,12	2,08	2,04	1,99	1,95	1,91	1,88	1,84	
24	4,26	3,40	3,01	2,78	2,62	2,51	2,42	2,36	2,30	2,25	2,22	2,18	2,15	2,13	2,11	2,03	1,98	1,94	1,89	1,84	1,80	1,77	1,73	
30	4,17	3,32	2,92	2,69	2,53	2,42	2,33	2,27	2,21	2,16	2,13	2,09	2,06	2,04	2,01	1,93	1,89	1,84	1,79	1,74	1,70	1,66	1,62	
40	4,08	3,23	2,84	2,61	2,45	2,34	2,25	2,18	2,12	2,08	2,04	2,00	1,97	1,95	1,92	1,84	1,79	1,74	1,69	1,64	1,59	1,55	1,51	
60	4,00	3,15	2,76	2,53	2,37	2,25	2,17	2,10	2,04	1,99	1,95	1,92	1,89	1,86	1,84	1,75	1,70	1,65	1,59	1,53	1,48	1,44	1,39	
100	3,94	3,09	2,70	2,46	2,31	2,19	2,10	2,03	1,97	1,93	1,89	1,85	1,82	1,79	1,77	1,68	1,63	1,57	1,52	1,45	1,39	1,34	1,28	
200	3,89	3,04	2,65	2,42	2,26	2,14	2,06	1,98	1,93	1,88	1,84	1,80	1,77	1,74	1,72	1,62	1,57	1,52	1,46	1,39	1,32	1,26	1,19	
+ ∞	3,84	3,00	2,61	2,37	2,21	2,10	2,01	1,94	1,88	1,83	1,79	1,75	1,72	1,69	1,67	1,57	1,52	1,46	1,39	1,32	1,24	1,17	1,01	

Logica dell'ANOVA: scomposizione della variabilità delle reazioni dei Ss (devianza) in un disegno classico a 1 fattore btw (o VI)

Per verificare statisticamente le ipotesi H_0 e H_1 ,
le devianze vengono divise per i rispettivi gradi di libertà
(GL, N-1 per DEV TOT, k -1 per DEV TRA, N - k per DEV ENTRO,
dove k = livelli di una VI)
e si ottengono nuovi valori noti come quadrati medi
(MS, mean squares) o varianze;
dal rapporto tra $MS_{TRA} / MS_{ENTRO} = F$
un rapporto che segue una distribuzione nota, detta appunto F

Cfr esercitazioni in excel

$$F = \frac{dev.regress/k}{dev.residua/(N - k - 1)}$$



$$F = MS_{TRA} / MS_{ENTRO}$$

$$MS_{TRA} = \sum \sum (M_{gruppo_j} - M_{GEN})^2 / k-1$$

$$MS_{entro} = \sum \sum (y_{ij} - M_{gruppo_j})^2 / N-k$$

Come quantificare l'effetto della VI dicotomica sulla VD, in un'ANOVA btw?

Eta quadro rivela la grandezza dell'effetto,
quanta parte di variabilità delle VD è spiegata dalla VI;
eta quadro corrisponde al **rapporto tra DEV_{TRA}/DEV_{TOTALE}**

sommatoria DEV TRA i gruppi	355.89
sommatoria DEV ENTRO i gruppi	9099.09
sommatoria dev totale	9454.99
F test	7.04
gl = 1, 180	
eta quadro	0.04

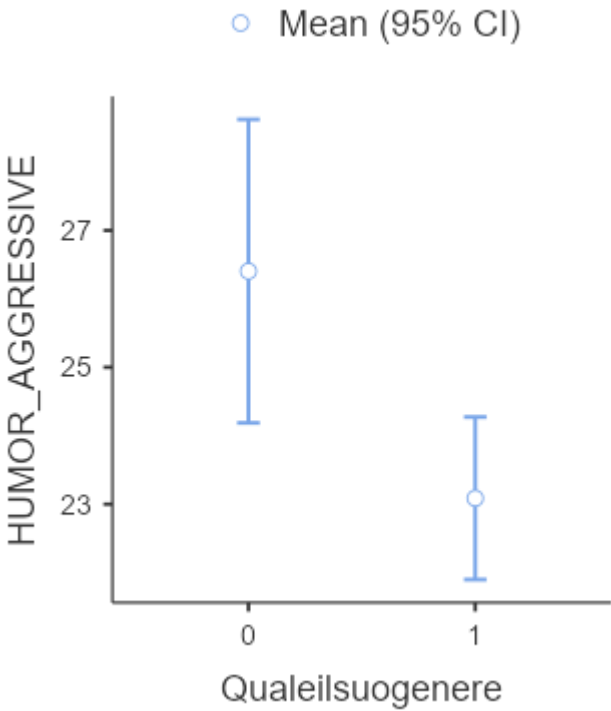
One-Way ANOVA

One-Way ANOVA (Welch's)

	F	df1	df2	p
HUMOR_AGGRESSIVE	7.03	1	67.4	0.010

Group Descriptives

	Qualeilsuogenere	N	Mean	SD	SE
HUMOR_AGGRESSIVE	0	42	26.4	7.12	1.098
	1	140	23.1	7.11	0.601



Linear Regression

Model Fit Measures

Model	R	R ²
1	0.194	0.0376

Metti a confronto con esito ANOVA 😊

Omnibus ANOVA Test

	Sum of Squares	df	Mean Square	F	p
Qualeilsuogenere	356	1	355.9	7.04	0.009
Residuals	9099	180	50.6		

Note. Type 3 sum of squares

Model Coefficients - HUMOR_AGGRESSIVE

Predictor	Estimate	SE	t	p	Stand. Estimate
Intercept	26.40	1.10	24.07	< .001	
Qualeilsuogenere	-3.32	1.25	-2.65	0.009	-0.194

Logica dell'ANOVA: scomposizione della variabilità delle reazioni dei Ss (devianza) in un disegno classico a 1 fattore btw (o VI)

*Come usare F per prendere decisioni circa le ipotesi?
Quando F indica che si può respingere H_0 ?*

F è tanto maggiore quanto la varianza TRA i gruppi è maggiore rispetto a quella ENTRO i gruppi: pertanto **quanto più F è elevato, tanto più esso sarà statisticamente significativo ($F > 0$)** ovvero tanto più le medie dei gruppi saranno distanti tra loro e tanto maggiore sarà la probabilità che siano statisticamente differenti;

altrimenti, se F è piccolo (cioè statisticamente $F = 0$), la variabilità tra i gruppi è sostanzialmente equivalente a quella entro i gruppi, vale a dire la variabilità legata al trattamento non differisce statisticamente da quella legata ad elementi residuali, alla varianza d'errore.

La decisione viene presa in base alla regola per cui **$p \leq \alpha$** .

Disegni sperimentali e a misure ripetute (non sperimentali) within con 1 VI

Il disegno sperimentale classico

Obiettivo: verificare l'ipotesi secondo cui la VI influenza causalmente le reazioni dei Ss alle prove che costituiscono il post-test (VD)

	Pre-test	Trattamento	Post-test	
Condiz speriment	si	si	si	→ Medie a confronto
Condiz contr	si	no	si	

Vantaggi:

- equivalenza dei gruppi rispetto a possibili var intervenienti
- con conseguente riduzione dell'impatto delle differenze individuali
- e massimizzazione degli eventuali effetti della VI

Disegni sperimentali e a misure ripetute (non sperimentali) within con 1 VI

Esempio di un disegno sperimentale

Le parti di un volto si riconoscono più facilmente se prese nel contesto di un volto anziché isolate? Ipotesi: SI.

VD: tempi di reazione (TR) nella decisione

VI: stimoli rappresentati da volti:

2 livelli: stimoli con contesto

stimoli senza contesto

Compito: decidere se due parti di un volto, presentate via computer, sono simili/diverse

Analisi: Confronto i TR osservati nelle due condizioni

Esempio di un disegno a misure ripetute non sperimentale:

Ipotesi: L'intelligenza fluida aumenta con l'età in età scolare

VD = punteggio alle Matrici di Raven

VI: tempo, a 2 livelli: I e II occasione di rilevazione a distanza di 1 anno

Target: Bambini di 7 anni

Compito: risolvere le matrici di Raven, stessi partecipanti nelle 2 occasioni di misurazione

Analisi: Confronto i punteggi osservati nei 2 momenti temporali

Disegni sperimentali e a misure ripetute (non sperimentali) within con 1 VI

In ogni misurazione, il punteggio osservato per un VD dipende da 3 elementi:

- effetto della VI
- effetto di caratteristiche proprie della persona/partecipante (componente sistematica)
- effetto del caso

Nei disegni MR o within è possibile stimare la componente dovuta alla persona che si assume rimanere relativamente costante da una misurazione all'altra, [se vicina nel tempo](#), e per questo può essere scorporata dalla componente di "errore" casuale (questa scomposizione invece non è possibile quando si ha una sola rilevazione per Ss come nei disegni between)

La componente sistematica influenza in modo costante la prestazione dei Ss attraverso i livelli della VI che pertanto saranno correlati, al di là degli effetti dovuti alla VI

Le differenze individuali non sono fonte di variabilità d'errore, perché costanti da una prova all'altra; l'errore è dovuto a variazioni accidentali da una risposta all'altra

Disegni sperimentali e a misure ripetute (non sperimentali) within con 1 VI

$$y_{ij} = M_{GEN} + (M_{prova_j} - M_{GEN}) + (y_{ij} - M_{prova_j})$$

Scomposizione della variabilità totale:

$$\sum \sum (y_{ij} - M_{GEN})^2 = \sum \sum (M_{prova_j} - M_{GEN})^2 + \sum \sum (y_{ij} - M_{prova_j})^2$$

dove la DEV_K è la DEV between o tra le prove

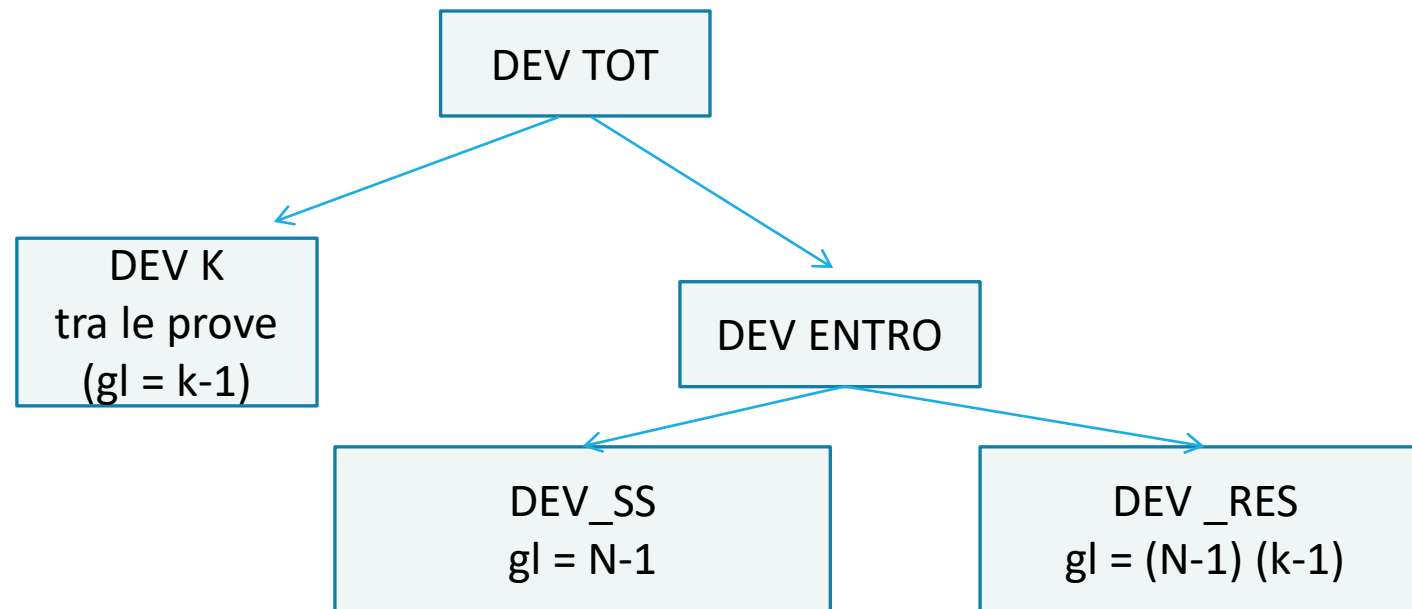
e la DEV_{WITHIN} si scompone in 2 parti:

$$DEV \text{ tra i Soggetti : } DEV_{SS} = \sum \sum (M_{yi} - M_{GEN})^2$$

$$DEV \text{ residua : } DEV_{RES} = DEV_{WITHIN} - DEV_{SS}$$

DEV_{SS} : le differenze individuali non sono fonte di variabilità d'errore, perché costanti da una prova all'altra; l'errore è dovuto a variazioni accidentali da una risposta all'altra (DEV_{RES})

Disegni sperimentali e a misure ripetute (non sperimentali) within con 1 VI



Nell'ANOVA per misure ripetute il rapporto F corrisponde a

$$F = (DEV_K / gl) / (DEV_{RES} / gl)$$

Disegni sperimentali e a misure ripetute (non sperimentali) within con 1 VI

Ipotesi: le parti di un volto si riconoscono più facilmente se prese nel contesto di un volto anziché isolate

Compito: decidere se due parti di un volto, presentate via computer, sono simili/diverse, VD quantificata con tempi di reazione (TR)

Condizioni: presentazione stimoli con contesto / senza contesto (VI), in ordine randomizzato per ogni soggetto (N = 10)

Risultati:

1. statistiche descrittive relative ai TR osservati nelle 2 condizioni

Estimates

Measure: MEASURE_1

factor1	Mean	Std. Error	95% Confidence Interval	
			Lower Bound	Upper Bound
CON	275,500	7,534	258,457	292,543
SENZA	289,500	7,403	272,753	306,247

1. Grand Mean

Measure: MEASURE_1

Mean	Std. Error	95% Confidence Interval	
		Lower Bound	Upper Bound
282,500	7,433	265,685	299,315

Disegni sperimentali e a misure ripetute (non sperimentali) within con 1 VI

I due valori medi osservati nelle due condizioni sono tra loro statisticamente diversi (H_1) o no (H_0) ?

Tests of Within-Subjects Contrasts

Measure: MEASURE_1

Source	factor1	Type I Sum of Squares	df	Mean Square	F	Sig.
factor1	Linear	980,000	1	980,000	91,875	,000
Error(factor1)	Linear	96,000	9	10,667		

*Il rapporto F osservato è statisticamente significativo;
pertanto la risposta è SI.*

Disegni sperimentali e a misure ripetute (non sperimentali) within con 1 VI

Quale la grandezza dell'effetto della VI sulla VD?

Tests of Within-Subjects Contrasts

Measure: MEASURE_1

Source	factor1	Type I Sum of Squares	df	Mean Square	F	Sig.
factor1	Linear	980,000	1	980,000	91,875	,000
Error(factor1)	Linear	96,000	9	10,667		

$$\eta^2 = DEV_K / DEV_K + DEV_{ERR}$$

$$\eta^2 = 980 / (980 + 96) = 0,91$$

NB nei disegni within è comune trovare un livello di varianza spiegata elevato poiché l'effetto del trattamento (la variabilità delle media delle condizioni intorno alla media generale) viene stimato rispetto alla variabilità dell'errore casuale solamente, non dall'errore generale che include anche la variabilità sistematica

Disegni sperimentali e a misure ripetute (non sperimentali) within con 1 VI

*Vi sono differenze **tra** i soggetti?*

Tests of Between-Subjects Effects

Measure: MEASURE_1

Transformed Variable: Average

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Intercept	1596125,000	1	1596125,000	1444,457	,000
Error	9945,000	9	1105,000		

*Il rapporto F osservato è statisticamente significativo;
pertanto la risposta è SI, vi sono differenze
nelle reazioni individuali dei partecipanti.*