

Statistica di base degli studi Real-life II

Scardapane Marco

Diapositiva preparata da MARCO SCARDAPANE e ceduta alla Società Italiana di Diabetologia.
Per ricevere la versione originale si prega di scrivere a siditalia@siditalia.it



Center for Outcomes Research and Clinical Epidemiology

Studi osservazionali vs RCTs / 1

- Una delle maggiori minacce per la validità di una stima, e.g. dell'effetto di un trattamento, è il selection bias;
- Nei trial clinici tale pericolo non sussiste, grazie all'assegnazione casuale in condizioni sperimentali controllate dei pazienti (**randomizzazione**);
- La probabilità di ricevere un trattamento è la stessa, e non dipende da fattori clinici o socio-demografici del paziente.

Diapositiva preparata da MARCO SCARDAPANE e ceduta alla Società Italiana di Diabetologia.
Per ricevere la versione originale si prega di scrivere a siditalia@siditalia.it

Studi osservazionali vs RCTs / 2

- D'altra parte, a volte non è ragionevole condurre un trial clinico, sia per ragioni di eticità che di fattibilità;
- Con l'avvento dei grandi database (real-world data, disease registries, etc.), gli studi osservazionali hanno a loro disposizione una ricchezza ed eterogeneità di informazioni che per natura un trial clinico non può mai dare;
- Con queste basi di dati è possibile esplorare l'efficacia di un trattamento in un ampio numero di sottogruppi di pazienti;
- E' possibile analizzare eventi rari, o effect size piccoli;
- E' possibile quindi ottenere una maggiore generalizzabilità (rispetto a un trial) ad un minor costo e tempo.

Diapositiva preparata da MARCO SCARDAPANE e ceduta alla Società Italiana di Diabetologia.
Per ricevere la versione originale si prega di scrivere a siditalia@siditalia.it

Una possibile soluzione ...

- Tutti questi vantaggi non hanno senso se non si tiene conto dei potenziali bias che insidiano qualunque studio osservazionale;
- Aggiustare in fase di disegno dello studio o in fase di analisi (analisi multivariata) spesso non è praticabile, per due motivi:
 - o Le variabili potenzialmente confondenti sono numerose;
 - o Non tutti i pazienti cadono nella totalità di valori potenzialmente assumibili da queste variabili (problema dell'estrapolazione).
- E' necessario trovare un modo per «mimare» la randomizzazione in condizioni assolutamente non controllate come quelle degli studi osservazionali.
- Nel 1983 Rosenbaum P. e Rubin D. idearono una misura sintetica in grado di «simulare» l'assegnazione casuale tipica dei trial negli studi osservazionali, il **propensity score**.

Il propensity score

- Il **propensity score (PS)** di un paziente è definito come la sua probabilità di ricevere il trattamento, condizionatamente a un insieme di covariate osservate;
- In un trial clinico con allocation ratio 1:1, il PS di un paziente è costante e pari a $1/2$;
- In uno studio osservazionale varia, e dipende dall'insieme di covariate;
- Il PS va quindi stimato, con una regressione logistica;
- In sostanza quindi la soluzione proposta è a due stadi: anziché aggiustare direttamente per i confounder...
 1. Si riassume l'informazione dei confounder in un indice sintetico, il PS;
 2. Si aggiusta l'effetto del trattamento per il PS.

Diapositiva preparata da MARCO SCARDAPONE e ceduta alla Società Italiana di Diabetologia.
Per ricevere la versione originale si prega di scrivere a siditalia@siditalia.it

Il propensity score - un esempio / 1

- Supponiamo di voler analizzare il rischio di sviluppare una patologia, e abbiamo un gruppo di trattati (T) e non trattati (NT):

	Events	Total (%)
T	60	200 (30.0%)
NT	579	800 (72.4%)

Unadj RR=0.39 (95% CI 0.34-0.46) p<0.0001

- Supponiamo anche che i due gruppi siano sbilanciati per età e sesso:

	T	NT	p-value
n	200	800	
Age	59.9±10.2	67.4±10.2	<0.0001
Sex (Male)	27.5	47.1	<0.0001

Il propensity score - un esempio / 2

- Costruiamo un modello logistico per il nostro PS, otterremo:

$$\text{logit}(PS) = -3.05 - 0.06 \times \text{Age} - 0.73 \times (\text{Sex} = \text{Male})$$

- Es. un uomo di 40 anni avrà: $\text{logit}(PS) = -3.05 - 0.06 \times 40 - 0.73 \times 1 = -6.18 \rightarrow PS = \frac{\exp(-6.18)}{1 + \exp(-6.18)} = 0.002$
- Es. una donna di 70 anni avrà: $\text{logit}(PS) = -3.05 - 0.06 \times 70 - 0.73 \times 0 = -7.25 \rightarrow PS = \frac{\exp(-7.25)}{1 + \exp(-7.25)} = 0.001$
- Aggiustiamo l'effetto del trattamento per il PS, ad es. tramite matching:

	T	Match NT	p-value
n	199	199	
Age	60.2 ± 10.3	60.0 ± 10.2	0.51
Sex (Male)	26.6	27.6	0.79

$$RR = 0.52 \text{ (95\% CI 0.41-0.63) } p < 0.0001$$

Come costruire il modello?

- Nelle situazioni reali i database hanno decine se non centinaia di variabili.
- Quante e quali inserire nel modello dunque?
- In generale bisogna inserire tutte le covariate necessarie e sufficienti per far sì che i gruppi siano bilanciati, anche le loro interazioni.
- Tradotto: non risparmiarsi ma né esagerare.
- In assenza di linee guida cliniche, è possibile usare metodi di selezione automatica (e.g. stepwise), in versione riadattata.

Diapositiva preparata da MARCO SCARDAPANE e ceduta alla Società Italiana di Diabetologia.
Per ricevere la versione originale si prega di scrivere a siditalia@siditalia.it

Come aggiustare per il PS?

- Una volta scelto il modello definitivo per il PS, è necessario scegliere anche come usare il PS stimato ai fini di aggiustamento;
- Nell'es. precedente ne abbiamo visto uno (matching), ma ne esistono altri!
- Qualche esempio:
 - o Adjusted regression con PS come variabile continua;
 - o Adjusted regression con PS in quantili;
 - o Analisi stratificata sul PS in quintili;
 - o Inverse probability weighting regression.
- Ognuno di questi metodi ha lo stesso scopo, ma dà risultati diversi in termini di bilanciamento ottenuto (i.e. bias) e di eventuale perdita di pazienti (i.e. precision).

Diapositiva preparata da MARCO SCARDAPANE e ceduta alla Società Italiana di Diabetologia.
Per ricevere la versione originale si prega di scrivere a siditalia@siditalia.it

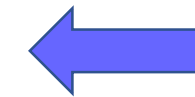
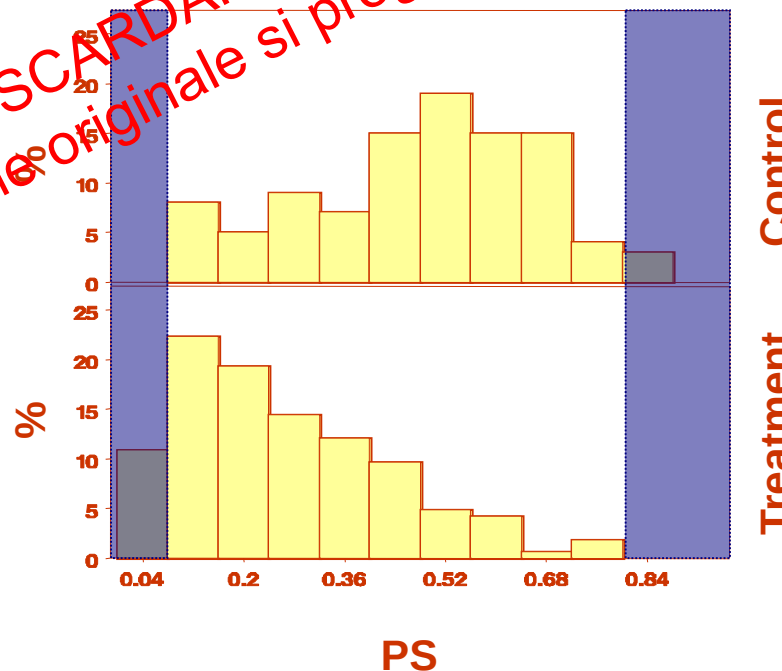
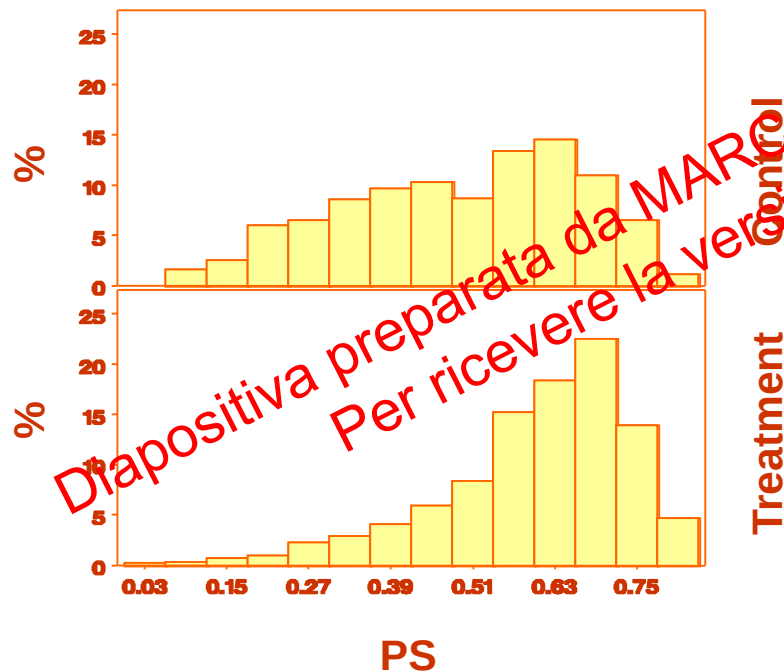
Bias vs precisione: l'eterna lotta

- In ogni ambito della statistica, e qui non c'è eccezione, c'è sempre un perenne tradeoff fra validità (o bias) e precisione;
- Se si minimizza l'uno lo si fa a scapito dell'altro;
- Un esempio è l'uso del matching: più rendiamo stringenti i criteri di matching, più le nostre coppie saranno simili (→ meno bias), meno pazienti riusciremo a matchare e quindi dovremo scartare → diminuzione del nostro campione → meno precisione;
- L'aspetto chiave collegato al dilemma di sopra è **l'overlapping**.

Diapositiva preparata da MARCO SCARDAPANE e ceduta alla Società Italiana di Diabetologia.
Per ricevere la versione originale si prega di scrivere a siditalia@siditalia.it

L'overlapping

- Se si usa il matching come metodo di aggiustamento, l'overlapping dei pazienti va sempre verificato;
- Quando non tutti i pazienti trattati si matchano ai rispettivi controlli per PS, si parla di partial overlapping;
- Es.



Partial overlapping

I limiti del PS

- Ricordate, la randomizzazione garantisce il bilanciamento nei due gruppi dei fattori osservati e **ignoti**;
- Il PS tiene conto solo dei fattori noti (quelli inclusi nel modello logistico);
- Che fare dunque? Uno strumento utile in queste situazioni è la **sensitivity analysis**.

Diapositiva preparata da MARCO SCARDAPANE e ceduta alla Società Italiana di Diabetologia.
Per ricevere la versione originale si prega di scrivere a siditalia@siditalia.it

La sensitivity analysis

- In sostanza la sensitivity analysis consta di 3 passi;
 1. Si simula l'introduzione di un generico confounder ignoto;
 2. Si simulano diversi scenari in termini di forza dell'effetto del confounder e di sbilanciamento della prevalenza del confounder nei due gruppi;
 3. Si calcola come varia la nostra stima dell'effetto del trattamento nei diversi scenari.
- *Es. il risultato di una PS-quintiles adjusted Cox è: $HR=0.19$ (95% CI 0.07-0.51; $p=0.001$)*

HR Confounder	Prev diff	Cox HR	95 % CI
3	0.6	0.38	0.14-1.02
4	0.5	0.41	0.15-1.10
5	0.4	0.41	0.15-1.10
6	0.3	0.38	0.14-1.02