

Potential Outcomes & Treatment Effect

An Observational Study is an empirical investigation in which "the objective is to elucidate cause & effect relationships" in which it is not feasible to use controlled experimentation, e.g. assigning people to smoking.

Consider a treatment Z , e.g. $Z \in K = \{T, C\}$
 T is treatment (e.g. smoking) C is control (e.g. non smoking)

For unit i , the Potential response τ_i is i 's outcome if he received T , e.g. $\tau_{Ti} = 1$ if i would get lung cancer had he/she smoked. Likewise, τ_{Ci} is i 's outcome if he received C .

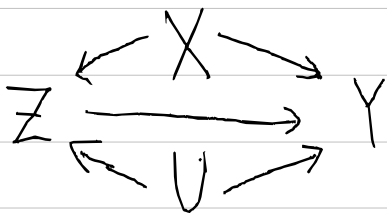
We are interested in $\tau_{Ti} - \tau_{Ci}$, the causal effect of T on the outcome. However, this is impossible, as only one can be observed. The unobserved outcome is called Counterfactual. So instead, we imagine selecting a patient at random, and taking expected value, to get the average treatment effect $ATE = E[\tau_T - \tau_C] = E[\tau_T] - E[\tau_C]$.

Causal Assumptions

Independence: The treatment actually assigned must be independent of the person's potential outcomes under treatment: $E[\tau_k] = E[\tau_k | Z = z] \quad \forall k, z \in K = \{T, C\}$

SUTVA: Stable Unit Treatment Value Assumption, this means τ_{iz} varies with z_i but not with any z_j , "No Interference between units"

No Unmeasured Confounders: To explain this assumption, we invoke the DAG (Directed Acyclic Graph), used to show causal relationships, where an arrow indicates causality and no arrow indicates independence. No unmeasured confounders says that if X are measured and may have an effect on Z and the observed outcome $Y = \tau_z$, then no unmeasured variable U exists such that:



e.g. No gene exists that makes you more likely to smoke and more likely to get lung cancer

Common ATE Tests

We wish to test the null hypothesis $H_0: \tau_T = \tau_C, \forall i$
Some common tests against this in different settings:

1. Fisher's Exact, if outcome is binary
2. Mantel-Haenszel, Outcome is binary but TWO or more strata. If the strata are matched pairs, it's called McNemar
3. t-test, paired or Unpaired, assume normal data
4. Wilcoxon's Rank Sum, Rank observed responses small to large, under no effect the sum of ranks of treated has known distribution, if $\tau_T > \tau_C$, sum should be large
5. Wilcoxon's Signed Rank, Used for paired data, rank absolute differences, sum ranks with $\tau_{Ti1} > \tau_{Ci2}$
6. Stratified Rank Sum, Wilcoxon for many strata
7. Aligned Rank, Superior to 6 for large # of strata
8. Median Test, Count # of treated to exceed within stratum median

Often, the exact p-values are difficult to calculate in large samples, so approximations (e.g. based on the normal using CLT) are often used

Treatment Effect Models

Additive treatment effect: The treatment raises or lowers the response by a fixed value ψ for all i . Obviously we cannot check it directly, but we can find the set of ψ 's such that some test of the null hypothesis

$H_0: \tau_{Ti} = \tau_{Ci} + \psi$ is not rejected. This is the confidence set for ψ under this model.

Dilated Treatment effects:

Multiplicative treatment effect: The treatment scales up or down the response by δ , the confidence set is generated from

$H_0: \tau_{Ti} = \delta \tau_{Ci}$, this is additive on log-scale:

$$H_0: \log(\tau_{Ti}) = \log(\tau_{Ci}) + \log(\delta)$$

Linear treatment effect: treat linear model of control:

$$H_0: \tau_{Ti} = \beta + \delta \tau_{Ci}$$

Some-Susceptible effect: Treat effect is 0 if $\tau_{Ci} < \tilde{\tau}$, but exists if $\tau_{Ci} > \tilde{\tau}$

Matching

Often, in observational data, there exist many measured confounders, that is $Z \xrightarrow{X_0} Y$, to determine the ATE using previously mentioned methods, we need to address these confounders, because Z is not assigned randomly, and we need to construct a "random" Z , Z that is independent of Y except through the causal relationship $Z \rightarrow Y$, by addressing these confounders.

One strategy is matching. The idea is that, assuming no unmeasured confounders, if treated person i 's covariates X_i are the same as control person j 's covariates X_j , then it is effectively "random" that person j was treated and person i was control instead of vice-versa.

A good matching balances the covariates between the matched treated group and matched controls. Often, this is assisted by a "propensity-score model".

A good matching is also often measured by standardized differences between matched T and C in covariate X :

$$d = \frac{\bar{X}_T - \bar{X}_C}{\sqrt{\frac{S_{X,T}^2 + S_{X,C}^2}{2}}} \quad \begin{array}{l} \text{Shifts} \\ \text{of } > 0.2 \\ \text{desired} \end{array}$$

Propensity Score & IPW

The Propensity Score $\pi(x_i) = P(Z_i = T | X = x_i)$ is the probability that i is treated given his covariates X_i . Often matching procedures are based on or involve a propensity score model of some kind, usually a logistic model: $\text{logit}(\pi(x_i)) = X\beta$, easy to fit

The IPW or inverse probability weighting estimator of $E[Y|Z]$ addresses X through $\pi(X)$, the idea is to downweight overrepresented (high probability of being observed) observations of Y and upweight underrepresented (low probability of being observed) observations of Y so that the resultant estimates of $E[Y_T]$, $E[Y_C]$ are unbiased. The resultant estimator of ATE is

$$\hat{\psi}_{IPW} = \frac{1}{n} \sum_{i=1}^n \left(\frac{Z_i Y_i}{\hat{\pi}(X_i)} - \frac{(1-Z_i) Y_i}{1 - \hat{\pi}(X_i)} \right) \quad \text{Since}$$

$$E[\hat{\psi}_{IPW}] = \frac{1}{n} \sum_{i=1}^n \frac{\pi(x_i) E[Y_i | Z_i=1, X_i]}{\pi(x_i)} - \frac{(1-\pi(x_i)) E[Y_i | Z_i=0, X_i]}{(1-\pi(x_i))} = E[Y | Z=1] - E[Y | Z=0]$$

assuming $E[\hat{\pi}(X_i)] = \pi(X_i)$, the propensity score model is good
Since X_i is given, $E[Z_i | X_i] = \pi(X_i)$

AIPW

The IPW can have high variance due to extreme weights, and requires $\pi(x_i)$ to be correctly specified. The AIPW involves an IPW and a second model, and AIPW is consistent if either is correctly specified, and is more robust than IPW.

Let $\hat{u}_1(x_i)$ be a model for Y_i based only on x_i , and $\hat{u}_0(x_i)$ be likewise, $\hat{u}_1$ for treated, $\hat{u}_0$ for control.

$$\hat{Y}_{AIPW} = \frac{1}{n} \sum_{i=1}^n \left[\frac{z_i(Y_i - \hat{u}_1(x_i))}{\hat{\pi}(x_i)} + \hat{u}_1(x_i) - \frac{(1-z_i)(Y_i - \hat{u}_0(x_i))}{\hat{\pi}(x_i)} - \hat{u}_0(x_i) \right]$$

Thus, $\hat{Y}_{AIPW}$ is $\hat{u}_1(x_i) - \hat{u}_0(x_i)$ except that the residuals $(Y_i - \hat{u}_1(x_i))$ and $(Y_i - \hat{u}_0(x_i))$ inverse weighted by the probability of being treated are added to the model. If both models are accurate, residuals cancel and we get $E[\hat{u}_1(x_i) - \hat{u}_0(x_i)] = ATE$, if only u is accurate the residuals will spread "evenly" up and down, since $E[Y_i - \hat{u}(x_i)]$ terms = 0 and result is consistent. Likewise, if only π is correct, the residual weights "fix" the bias in u . e.g., if $\hat{u}_1(x_i) = c_1$, $\hat{u}_0 = c_0$, AIPW can be shown to be just IPW.

Sensitivity Analysis

Often, we are concerned that perhaps there exists a hidden bias, an unmeasured confounder U , that affects both the treatment and the outcome, violating our assumptions. This assumption is untestable, but there are many ways to address the concern. One way is Sensitivity Analysis, which tells us how strong a hidden bias would have to be in order to change our conclusion.

For example: Concerned about gene X that both causes people to smoke and get lung cancer. Sensitivity analysis says that for our test to lose statistical significance the gene would have to be 4 times more present in smokers than nonsmokers.

Let j and k be units with same covariates X , one T one C , where $\pi_j \neq \pi_k$ due to a hidden bias U . I imagine we knew the odds ratio for units with matching X is at most Γ . When the test being used is a ranked sum, it can be shown that

Sensitivity Analysis

$$\frac{1}{\pi} \leq \frac{\pi_j(1-\pi_k)}{\pi_k(1-\pi_j)} \leq \pi \quad \forall_{j,k \text{ s.t. } X_j = X_k} \Rightarrow \frac{1}{1+\pi} \leq P(Z_i=1 | X_i \text{ and } X_j) \leq \frac{\pi}{1+\pi}$$

The change in π leads to a change in the distribution of the rank-sum null distribution. Suppose we are working with paired data, again T is a ranked sum. Let S be a pair, C_{S1} and C_{S2} be the treatment status of units in the pair, $d_S = Y_{S1}^T - Y_{S2}^C$, define

$$P_S^+ = \begin{cases} 0 & C_{S1} = C_{S2} = 0 \\ 1 & C_{S1} = C_{S2} = 1 \\ \frac{\pi}{1+\pi} & C_{S1} \neq C_{S2} \end{cases} \quad P_S^- = \begin{cases} 0 & C_{S1} = C_{S2} = 0 \\ 1 & C_{S1} = C_{S2} = 1 \\ \frac{1}{1+\pi} & C_{S1} \neq C_{S2} \end{cases}$$

min/max prob that the pair contributes to the ranked sum. Define T^+ as sum of S random variables where j th r.v. takes on value d_S with probability P_S^+ , and T^- as same except j takes on d_S with probability P_S^- . Mean and Variance are easily calculated, and can often be normal approximated (R_i 's ranks of d_i 's):

$$E[T^+] = \sum_{S=1}^n P_S^+ R_S \quad \xrightarrow{\text{if no S covariant}} \sum_{S=1}^n P_S^+ \frac{\pi}{1+\pi}$$

$$\text{Var}(T^+) = \sum_{S=1}^n P_S^+ (1-P_S^+) R_S^2 \quad \text{Replace } P_S^+ \text{ w/ } P_S^- \text{ to get result for } T^-$$

Sensitivity Interval

A Sensitivity Interval of Γ for a parameter β is the set of β 's for which the null hypotheses for which (say $H_0: \tau_{\tau_1} = \tau_{\tau_0} + \beta$ in additive model) cannot be rejected even with sensitivity parameter Γ (a hidden bias could be Γ times more present or $\frac{1}{\Gamma}$ times as present in treatment compared to control) working against the null hypothesis

Analogue of Confidence interval when Γ known to have some maximum value. For example:

$$SI_{\beta} : \left(\sup_{\beta} \left(\frac{T_{\beta, obs} - E[T_{\beta}^-]}{\sqrt{\text{Var}(T_{\beta, n})}} \right) \leq 1.96 \right), \inf_{\beta} \left(\frac{T_{\beta, obs} - E[T_{\beta, n}^+]}{\sqrt{\text{Var}(T_{\beta, n})}} \geq -1.96 \right)$$

Using normal approx. for Wilcoxon RS Sensitivity Analysis.

Nested Hypotheses

Another method for testing while accounting for a specific hidden bias is to use Nested Hypotheses. First, test the treatment group against one control group that you can assume is not subject to one kind of hidden bias. Then, only if the test is significant, test against a different control group that has a different potential hidden bias. If this is also significant, draw the conclusion.

Example: Miners of lead/zinc in a town might be subject to genetic damage from radon exposure. Use two control groups: 1) Non-miners who live near the mine 2) matched/related field workers distant from the mine. The hidden bias controlled by 1) is effect of lead pollution on local environment. The hidden bias controlled by 2) is that of difference in occupation/class that could contribute to increased genetic damage.

Difference in Differences

Used to estimate the effect of an intervention by comparing the difference between the treated group and control group before and after the intervention. This is often used to evaluate a sudden policy change.

Compute the treatment effect as

$$\delta = (Y_{T, \text{post}} - Y_{T, \text{pre}}) - (Y_{C, \text{post}} - Y_{C, \text{pre}})$$

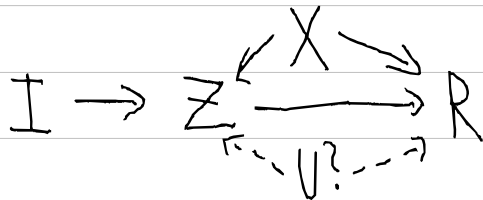
Example: Philadelphia babies, between 1997 and 2007 many hospitals closed obstetric units (12 of 19) with no similar reduction across U.S., did this lead to more newborn deaths? Matched Philadelphia is constructed for Y_C using similar neighborhoods from across U.S.

Again, Y_{it} 's should be matched to Y_{ic} 's so that $\hat{\delta} = (\bar{Y}_{T, \text{post}} - \bar{Y}_{T, \text{pre}}) - (\bar{Y}_{C, \text{post}} - \bar{Y}_{C, \text{pre}})$ is not affected by measured confounders

Instrumental Variables

Often, we cannot control the treatment of interest and suspect that there is unmeasured confounding of treatment, but there is another variable I associated with the treatment we know to be independent of Z . If I can be randomized, we call this a randomized encouragement design.

Ex: Randomly encourage some mothers not to smoke, cannot assign to smoking or non-smoke but can assign. encouragement against smoking, which will affect treatment and enc. against smoking can be assigned randomly. Let R be baby weight (outcome), $Z = 0, 1$, non-smoke or smoke, $X =$ covariates (e.g. mom age), $I = 1, 0$, encourage or don't



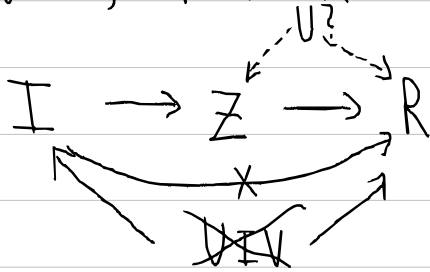
We call I an instrumental variable (IV) and can use IV to eliminate the effect of U

Instrumental Variables

Assumptions needed for a valid IV:

1. Exclusion Restriction: I affects R only through Z

2. Unmeasured Confounding of IV: I has no unmeasured confounders with R



Note in 3. that $Z_i(1)$ is i 's potential treatment if encouraged, $Z_i(0)$ is i 's potential treatment if not encouraged

Sometimes the I.V. is not from an encouragement design, instead some other naturally constructed "random encouragement" can be used as an I.V., e.g. living near 2-year or 4-year college

3. NO "defiers" whose potential outcomes are flipped, i.e. encouragement causes them to not be in treated group when they otherwise would be, $Z(1)=0$, $Z(0)=1$.

Other terminology: compliers have $Z(1)=1$, $Z(0)=0$, always takers have $Z(1)=1$, $Z(0)=1$, never takers have $Z(1)=0$, $Z(0)=0$

2SLS

From our IV analysis, we wish to estimate $\beta = E[R_i | Z_i=1, Z_i(1)=1, Z_i(0)=0] - E[R_i | Z_i=0, Z_i(1)=1, Z_i(0)=0]$

The effect of treatment on compliers. Since we do not have counterfactuals for always/never takers, no conclusions can be drawn about treatment effect for them.

Two-Stage Least Squares is an unbiased estimate of β , (if component models are correct) and first estimates $E[Z_i | X_i, I_i] = \alpha$ and then uses $\hat{\alpha}$ as a predictor of R_i , and create a another model with $X_i, \hat{\alpha}$,

Stage 1: $Z_i = \pi_0 + \pi_1 I_i + \pi_2 X_i + \nu_i$ for example
 $\hat{\alpha}_i = \hat{E}[Z_i | I_i, X_i]$ from this model

Stage 2: $R_i = \beta_0 + \beta_1 \hat{\alpha}_i + \beta_2 X_i + \epsilon_i$ where
our FVnd $\hat{\beta}_1$ is the estimator of interest

$$\hat{\beta}_{\text{TSLs}} = \frac{\hat{E}[R_i | I_i=1] - \hat{E}[R_i | I_i=0]}{\hat{E}[Z_i | I_i=1] - \hat{E}[Z_i | I_i=0]} \rightarrow \frac{P(Z_i(0)=0, Z_i(1)=1) \beta}{P(Z_i(0)=0, Z_i(1)=1)}$$

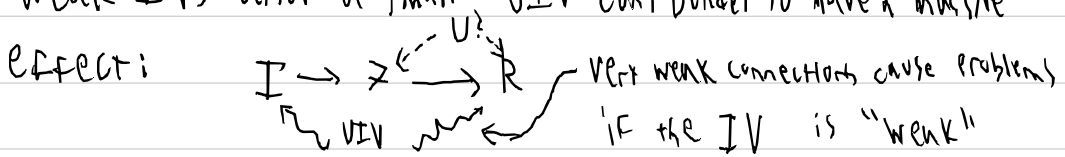
$= \beta$, as before tests exist to form CI's for β

Instrumental Variables

The prevalence ratio of a binary characteristic X_i among the compliers compared to the full population is $P(X_i=1 | Z_i(1)=1, Z_i(0)=0) / P(X_i=1)$, ratio of the prevalence of the trait in the "complier" group compared to the whole population. This helps tell us for what types of units are the IV results legit.

Assuming monotonicity,
$$\frac{E[Z_i | I_i=1, X_i=1] - E[Z_i | I_i=0, X_i=1]}{E[Z_i | I_i=1] - E[Z_i | I_i=0]}$$
 the prev ratio is

We can "guess" whether the IV is independent of unmeasured confounders by checking if it has a strong relationship with measured confounders. However, the IV is independent of everything. We also want to check the "strength" of the IV, this is done through variation of compliers. Ideally, we would like $P(\text{complier}) \approx 0.4$, a $P(\text{complier}) < 0.2$ is very weak. Weak IVs allow a small UIV confounder to have a massive effect:



Slightly biased strong IV generally better than a weak IV