

Causal inference with confounders missing not at random

By S. YANG

*Department of Statistics, North Carolina State University, 2311 Stinson Drive, Raleigh,
North Carolina 27695, U.S.A.*

syang24@ncsu.edu

L. WANG

*Department of Statistical Sciences, University of Toronto, 100 St. George Street, Toronto,
Ontario M5S 3G3, Canada*

linbo.wang@utoronto.ca

AND P. DING

*Department of Statistics, University of California, 367 Evans Hall, Berkeley,
California 94720, U.S.A.*

pengdingpku@berkeley.edu

SUMMARY

It is important to draw causal inference from observational studies, but this becomes challenging if the confounders have missing values. Generally, causal effects are not identifiable if the confounders are missing not at random. In this article we propose a novel framework for non-parametric identification of causal effects with confounders subject to an outcome-independent missingness, which means that the missing data mechanism is independent of the outcome, given the treatment and possibly missing confounders. We then propose a nonparametric two-stage least squares estimator and a parametric estimator for causal effects.

Some key words: Completeness; Identifiability; Ill-posed inverse problem; Integral equation; Outcome-independent missingness; Two-stage least squares estimator.

1. INTRODUCTION

Causal inference plays an important role in biomedical studies and social sciences. If all the confounders of the treatment-outcome relationship are observed, one can use standard techniques, such as propensity score matching, subclassification and weighting, to adjust for confounding (e.g., [Rosenbaum & Rubin, 1983](#); [Imbens & Rubin, 2015](#)).

Much less work has been done on the case where confounders have missing values. [Rosenbaum & Rubin \(1984\)](#) and [D'Agostino Jr & Rubin \(2000\)](#) developed a generalized propensity score approach. Under a modified unconfoundedness assumption, they showed that adjusting for the missing pattern and the observed values of confounders removes all confounding bias, and hence the causal effects are identifiable. Moreover, the balancing property of the propensity score carries over to the generalized propensity score. Standard propensity score methods can therefore be used to estimate the causal effects. However, the modified unconfoundedness assumption implies that units may have different confounders depending on the missing pattern, which is often difficult

to justify scientifically. An alternative approach assumes that the confounders are missing at random (Rubin, 1976). Under this assumption, both the full data distribution and the causal effects are identifiable, and multiple imputation can be used to obtain estimates of the causal effects (Rubin, 1987; Qu & Lipkovich, 2009; Crowe et al., 2010; Mitra & Reiter, 2011; Seaman & White, 2014). In practice, however, the missing pattern often depends on the missing values themselves, a scenario commonly known as missing not at random (Rubin, 1976). Multiple-imputation methods may fail to provide valid inference in this scenario. See Mattei (2009) for a comparison of various methods and Lu & Ashmead (2018) for a sensitivity analysis.

Causal inference with confounders missing not at random is challenging because neither the full data distribution nor the causal effects are identifiable without further assumptions. We consider a novel setting in which the confounders are subject to an outcome-independent missingness; that is, the missing data mechanism is independent of the outcome, given the treatment and possibly missing confounders. This outcome-independent missingness is plausible if the outcome happens after the covariate measurements and missing data indicators. To identify the causal effects in this setting, we formulate the identification problem as solving an integral equation, and show that the identification of the full data distribution is equivalent to the existence of a unique solution to an inverse problem. This new perspective allows us to establish a general condition for identifiability of the causal effects. Our condition generalizes existing results for discrete covariates and outcome (Ding & Geng, 2014). Motivated by the identification result, we develop a nonparametric two-stage least squares estimator by solving the sample analogue of the integral equation. To avoid the curse of dimensionality, we further develop parametric likelihood-based methods.

2. SET-UP AND ASSUMPTIONS

2.1. Potential outcomes, causal effects and unconfoundedness

We use potential outcomes to define causal effects (Neyman, 1923; Rubin, 1974). Suppose that the binary treatment is $A \in \{0, 1\}$, with 0 and 1 being the labels for the control and active treatments, respectively. Each level of treatment a corresponds to a possibly multi-dimensional potential outcome $Y(a)$, representing the outcome had the subject, possibly contrary to the fact, been given treatment a . The observed outcome is $Y = Y(A) = AY(1) + (1 - A)Y(0)$. Let $X = (X_1, \dots, X_p)$ be a vector of p -dimensional pre-treatment covariates. We assume that a sample of size n consists of independent and identically distributed draws from the distribution of $\{A, X, Y(0), Y(1)\}$. The covariate-specific causal effect is $\tau(X) = E\{Y(1) - Y(0) \mid X\}$, and the average causal effect is $\tau = E\{Y(1) - Y(0)\} = E\{\tau(X)\}$. We focus on τ ; a similar discussion applies to the average causal effect on the treated, $\tau_{\text{ATT}} = E\{Y(1) - Y(0) \mid A = 1\} = E\{\tau(X) \mid A = 1\}$. The following assumptions are standard in causal inference with observational studies (Rosenbaum & Rubin, 1983).

Assumption 1. We have that $\{Y(0), Y(1)\} \perp\!\!\!\perp A \mid X$.

Assumption 2. There exist constants c_1 and c_2 such that $0 < c_1 \leq e(X) \leq c_2 < 1$ almost surely, where $e(X) = \text{pr}(A = 1 \mid X)$ is the propensity score.

Under Assumptions 1 and 2, $\tau = E\{E(Y \mid A = 1, X) - E(Y \mid A = 0, X)\}$ is identifiable from the joint distribution of the observed data (A, X, Y) . Rosenbaum & Rubin (1983) showed that $\{Y(0), Y(1)\} \perp\!\!\!\perp A \mid e(X)$, so adjusting for the propensity score removes all confounding. We can estimate τ through propensity score matching, subclassification or weighting.

2.2. Confounders with missing values and the generalized propensity score

We consider the case where X contains missing values. Let $R = (R_1, \dots, R_p)$ be the vector of missing indicators such that $R_j = 1$ if the j th component X_j is observed and 0 if X_j is missing. Let $\mathcal{R}$ be a subset of all possible values of R . We use 1_p to denote the p -vector of 1s and 0_p the p -vector of 0s. The missingness pattern $R = r \in \mathcal{R}$ partitions the covariates X into X_r and $X_{\bar{r}}$, the observed and missing parts of X , respectively. Using the standard notation, $X_R = X_{\text{obs}}$ and $X_{\bar{R}} = X_{\text{mis}}$ are the realized observed and missing covariates, respectively. For example, if $R_1 = 1$ and $R_j = 0$ for $j = 2, \dots, p$, then $X_R = X_1$ and $X_{\bar{R}} = (X_2, \dots, X_p)$. Assume that the full data are independent and identically distributed draws from $\{A, X, Y(0), Y(1), R\}$, and so the observed data are independent and identically distributed draws from (A, R, X_R, Y) . [Rosenbaum & Rubin \(1984\)](#) introduced the following modified unconfoundedness assumption.

Assumption 3. We have that $\{Y(0), Y(1)\} \perp\!\!\!\perp A \mid (X_R, R)$.

Under Assumption 3, the generalized propensity score $e(X_R, R) = \text{pr}(A = 1 \mid X_R, R)$ plays the same role as the usual propensity score $e(X) = \text{pr}(A = 1 \mid X)$ in the settings without missing covariates. [Rosenbaum & Rubin \(1984\)](#) showed that adjusting for $e(X_R, R)$ balances (X_R, R) and removes all confounding on average. Their approach has the advantage of requiring no assumptions on the missing data mechanism of X for the identification of causal effects. However, Assumption 3 implies that a pre-treatment covariate can be a confounder when it is observed, but is not a confounder when it is missing; this is often hard to justify scientifically. Moreover, if the covariate measurement occurs after the treatment assignment, then R is a post-treatment variable affected by A . In this case, even if A is completely randomized, Assumption 3 is unlikely to hold when conditioning on the post-treatment variable R ([Frangakis & Rubin, 2002](#)).

2.3. Missing data mechanisms of the confounders

Without Assumption 3, identification of causal effects relies on alternative assumptions on the missing data mechanism. We now describe existing approaches under different missingness mechanisms of the confounders, the first of which is missing completely at random ([Rubin, 1976](#)).

Assumption 4 (Missing completely at random). We have that $R \perp\!\!\!\perp (A, X, Y)$.

Assumption 4 requires that the missingness of confounders be independent of all variables (A, X, Y) . It implies $\tau = E\{\tau(X) \mid R = 1_p\}$ and thus justifies the complete-case analysis that uses only the units with fully observed confounders. This complete-case analysis is, however, inefficient as it discards all units with missing confounders. Moreover, confounders are rarely missing completely at random.

The second missingness mechanism is missing at random ([Rubin, 1976](#)).

Assumption 5 (Missing at random). We have that $R \perp\!\!\!\perp X \mid (A, Y)$.

Under Assumption 5, conditioning on the treatment and outcome, the missing mechanism of confounders is independent of the missing values themselves. Assumption 5 implies $f(A, X, Y) = f(A, Y)f(X \mid A, Y, R = 1_p)$, and therefore the joint distribution $f(A, X, Y)$ and its functionals, including τ , are all identifiable. [Rubin \(1976\)](#) showed that the missing data mechanism can be ignored in the likelihood-based and Bayesian inferences under Assumption 5. In this case, multiple imputation is a popular tool for causal inference (e.g., [Qu & Lipkovich, 2009](#); [Crowe et al., 2010](#); [Mittra & Reiter, 2011](#); [Seaman & White, 2014](#)).

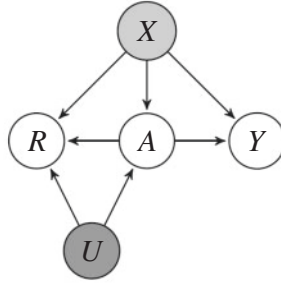


Fig. 1. A direct acyclic graph illustrating Assumptions 1 and 6; white nodes represent observed variables, the light grey node represents the variable with missing values, and the dark node represents an unmeasured variable U .

However, imputing the missing confounders based on $f(X_{\bar{R}} | X_R, A, Y) \propto f(X)f(A | X)f(Y | A, X)$ involves an outcome model in general (U.S. Department of Education, 2017), which is contrary to the suggestion of Rubin (2007) that the outcome should not be used in the design of an observational study. More importantly, missing at random is not plausible if the missing pattern depends on the missing values themselves. Instead, we consider the following missing data mechanism.

Assumption 6 (Outcome-independent missingness). We have that $R \perp\!\!\!\perp Y | (A, X)$.

Assumption 6 is plausible for prospective observational studies with covariates measured long before the outcome takes place (e.g., Hsu & Small, 2013; Hanna-Attisha et al., 2016). Figure 1 is a special causal diagram (Pearl, 1995) illustrating Assumptions 1 and 6. Graphically, A and Y have no common parents except for X , encoding Assumption 1, and R and Y have no common parents except A and X , encoding Assumption 6. Our framework allows for unmeasured common causes of R and A , as well as the dependence of R on the missing confounders $X_{\bar{R}}$. Moreover, it allows R to be a post-treatment variable affected by A . We give more graphical illustrations of Assumption 6 in the Supplementary Material.

We also make the following assumption to rule out degeneracy of the missing data mechanism.

Assumption 7. We have that $\text{pr}(R = 1_p | A, X, Y) > c_3 > 0$ almost surely for some constant c_3 .

3. NONPARAMETRIC IDENTIFICATION

3.1. Identification strategy

Assume that the distribution of (A, X, Y, R) is absolutely continuous with respect to some measure, with $f(A, X, Y, R)$ being the density or probability mass function. Under Assumptions 1 and 2, the key is to identify the joint distribution of $f(A, X, Y)$ because τ is its functional. The following identity relates the full data distribution to the observed data distribution:

$$f(A, X, Y, R = 1_p) = f(A, X, Y) \text{pr}(R = 1_p | A, X, Y). \quad (1)$$

The left-hand side of (1) is identifiable under Assumption 7. Therefore, the identification of $f(A, X, Y)$ relies on the identification of $\text{pr}(R = 1_p | A, X, Y)$. We now discuss how to identify $\text{pr}(R = 1_p | A, X, Y) = \text{pr}(R = 1_p | A, X)$ under Assumption 6.

3.2. Integral equation representation

Under Assumption 6, let

$$\xi_{ra}(X) = \frac{\text{pr}(R = r \mid A = a, X, Y)}{\text{pr}(R = 1_p \mid A = a, X, Y)} = \frac{\text{pr}(R = r \mid A = a, X)}{\text{pr}(R = 1_p \mid A = a, X)} \quad (a = 0, 1; r \in \mathcal{R}).$$

It then suffices to identify $\xi_{ra}(X)$, because it determines the missing data mechanism via

$$\text{pr}(R = r \mid A = a, X, Y) = \frac{\text{pr}(R = r \mid A = a, X, Y)}{\sum_{r' \in \mathcal{R}} \text{pr}(R = r' \mid A = a, X, Y)} = \frac{\xi_{ra}(X)}{\sum_{r' \in \mathcal{R}} \xi_{r'a}(X)}. \quad (2)$$

The following theorem shows that $\xi_{ra}(X)$ is a key term connecting the observed data distribution $f(A, X_r, Y, R = r)$ and the complete-case distribution $f(A, X, Y, R = 1_p)$. Throughout the paper, $\nu(\cdot)$ denotes a generic measure, such as the Lebesgue measure for a continuous variable or the counting measure for a discrete variable.

THEOREM 1. *Under Assumption 6, for any r and a , the following integral equation holds:*

$$f(A = a, X_r, Y, R = r) = \int \xi_{ra}(X) f(A = a, X, Y, R = 1_p) d\nu(X_{\bar{r}}). \quad (3)$$

Proof. The result follows because the observed data distribution is the complete-data distribution averaged over the missing data:

$$\begin{aligned} f(A = a, X_r, Y, R = r) &= \int f(A = a, X, Y, R = r) d\nu(X_{\bar{r}}) \\ &= \int \frac{\text{pr}(R = r \mid A = a, X, Y)}{\text{pr}(R = 1_p \mid A = a, X, Y)} f(A = a, X, Y, R = 1_p) d\nu(X_{\bar{r}}) \\ &= \int \xi_{ra}(X) f(A = a, X, Y, R = 1_p) d\nu(X_{\bar{r}}). \quad \square \end{aligned}$$

Theorem 1 is the basis of our identification analysis. In (3), $f(A = a, X_r, Y, R = r)$ and $f(A = a, X, Y, R = 1_p)$ are identifiable from the observed data. We have thus turned the identification of $\xi_{ra}(X)$ into the problem of solving for $\xi_{ra}(X)$ from (3). This requires additional technical assumptions, given below.

3.3. Bounded completeness and identification of the joint distribution

To motivate our identification conditions, we first consider the case of discrete X and Y , so that (3) becomes a linear system. To solve for $\xi_{ra}(X)$ from (3), we need the linear system to be nondegenerate.

PROPOSITION 1. *Under Assumption 6, suppose that X and Y are discrete, with $X_j \in \{x_{j1}, \dots, x_{jJ_j}\}$ for $j = 1, \dots, p$ and $Y \in \{y_1, \dots, y_K\}$. Let $q = J_1 \times \dots \times J_p$, and let Θ_a be a $K \times q$ matrix with the k th row being $f(X, y_k, R = 1_p, A = a)$ evaluated at all possible values of X . The distribution of (A, X, Y, R) is identifiable if $\text{Rank}(\Theta_a) = q$ for $a = 0, 1$.*

We relegate the proof to the Supplementary Material. For the special case of a binary X and a discrete Y , the rank condition in Proposition 1 is equivalent to $X \not\perp\!\!\!\perp Y \mid (A = a, R = 1)$ for $a = 0$ and 1, which is testable based on the observed data (Ding & Geng, 2014). For general cases, we

need to extend the rank condition that ensures the unique existence of $\xi_{ra}(X)$. We use the notion of bounded completeness for general X and Y , which is related to the concept of a complete statistic (Lehmann & Scheffé, 1950; Newey & Powell, 2003). Below, we say that a function $g(x)$ is bounded in $\mathcal{L}_1$ -metric if $\sup_x |g(x)| \leq c$ for some $0 < c < \infty$.

DEFINITION 1. A function $f(X, Y)$ is bounded complete in Y if $\int g(X)f(X, Y) d\nu(X) = 0$ implies $g(X) = 0$ almost surely for any measurable function $g(X)$ bounded in $\mathcal{L}_1$ -metric.

D'Haultfoeuille (2011) gave sufficient conditions for bounded completeness. Bounded completeness has also appeared in other identification analyses, such as nonparametric instrumental variable regression models (Darolles et al., 2011) and measurement error models (An & Hu, 2012).

We invoke the following assumption motivated by Theorem 1 and Definition 1.

Assumption 8. The joint distribution $f(A = a, X, Y, R = 1_p)$ is bounded complete in Y for $a = 0, 1$.

Remark 1. When X and Y are discrete with finite supports, Assumption 8 is equivalent to the rank condition in Proposition 1. For continuous X and Y , Assumption 8 requires that the dimension of Y be at least as large as the dimension of X in general. Moreover, Assumption 8 implies Assumption 2. We give more details for these results in the Supplementary Material.

Under Assumption 7, Assumption 8 is sufficient to ensure the existence and uniqueness of $\xi_{ra}(X)$ from (3). We state the result in the following theorem.

THEOREM 2. Under Assumptions 6–8, the distribution of (A, X, Y, R) is identifiable.

Proof. Suppose that $\xi_{ra}^{(1)}(X)$ and $\xi_{ra}^{(2)}(X)$ are two solutions to (3):

$$f(A = a, X_r, Y, R = r) = \int \xi_{ra}^{(k)}(X) f(A = a, X, Y, R = 1_p) d\nu(X_{\bar{r}}) \quad (k = 1, 2),$$

implying that $\int \{\xi_{ra}^{(1)}(X) - \xi_{ra}^{(2)}(X)\} f(A = a, X, Y, R = 1_p) d\nu(X_{\bar{r}}) = 0$. Integrating this identity with respect to X_r gives

$$\int \{\xi_{ra}^{(1)}(X) - \xi_{ra}^{(2)}(X)\} f(A = a, X, Y, R = 1_p) d\nu(X) = 0.$$

Assumption 7 implies that $\xi_{ra}(X)$ is bounded in $\mathcal{L}_1$ -metric, which further implies that $\xi_{ra}^{(1)}(X) - \xi_{ra}^{(2)}(X)$ is bounded in $\mathcal{L}_1$ -metric. Under Assumption 8, Definition 1 implies that $\xi_{ra}^{(1)}(X) - \xi_{ra}^{(2)}(X) = 0$ almost surely. Therefore, (3) has a unique solution $\xi_{ra}(X)$. Based on the definition of $\xi_{ra}(X)$, we can identify $\text{pr}(R = 1_p | A, X, Y)$ by (2). Finally, we identify $f(A, X, Y)$ through (1) as $f(A, X, Y) = f(A, X, Y, R = 1_p) / \text{pr}(R = 1_p | A, X, Y)$. $\square$

If the distribution of (A, X, Y) is identifiable, we can use a standard argument to show that τ and τ_{ATT} are identifiable under Assumption 1. In the next subsection we give explicit identification formulas for τ and τ_{ATT} , which form the basis for constructing the nonparametric estimator.

3.4. Nonparametric identification formulas for average causal effects

Under Assumptions 1 and 6–8, we can identify τ and τ_{ATT} in two steps. First,

$$\tau(X) = E(Y | A = 1, X) - E(Y | A = 0, X) \quad (4)$$

$$= E(Y | A = 1, X, R = 1_p) - E(Y | A = 0, X, R = 1_p), \quad (5)$$

where (4) follows from Assumption 1 and (5) follows from Assumption 6. Therefore, we can identify $\tau(X)$ using a complete-case analysis based on (5).

Second, under Assumptions 6–8, Theorem 2 shows that the distribution of (A, X, Y, R) is identifiable, which implies that the marginal distribution of X , $f(X)$, and the conditional distribution of X , $f(X | A = 1)$, are also identifiable. Therefore, both $\tau = E\{\tau(X)\}$ and $\tau_{\text{ATT}} = E\{\tau(X) | A = 1\}$ are identifiable. The following theorem summarizes these results and gives the explicit formulas.

THEOREM 3. *Under Assumptions 1 and 6–8, the average causal effect τ is identified by*

$$\tau = \sum_{a=0}^1 \int \tau(X) \frac{f(A = a, X, R = 1_p)}{\text{pr}(R = 1_p | A = a, X)} d\nu(X), \quad (6)$$

and the average treatment effect on the treated, τ_{ATT} , is identified by

$$\tau_{\text{ATT}} = \int \tau(X) \frac{f(X, R = 1_p | A = 1)}{\text{pr}(R = 1_p | A = 1, X)} d\nu(X), \quad (7)$$

where $\tau(X)$ is identified by (5), $\text{pr}(A = a, R = 1_p)$ and $f(A = a, X, R = 1_p)$ depend only on the observed data, and $\text{pr}(R = 1_p | A = a, X)$ can be identified from (2) and (3) for $a = 0, 1$.

Proof. First, we can identify the conditional distribution of X given $A = a$ by

$$f(X | A = a) = \frac{f(X, R = 1_p | A = a)}{\text{pr}(R = 1_p | A = a, X)} \quad (a = 0, 1).$$

Averaging $\tau(X)$ over $f(X | A = 1)$ yields the identification formula (7).

Second, we can identify the marginal distribution of X by

$$f(X) = \sum_{a=0}^1 f(A = a, X) = \sum_{a=0}^1 \frac{f(A = a, X, R = 1_p)}{\text{pr}(R = 1_p | A = a, X)}.$$

Averaging $\tau(X)$ over the above distribution gives the identification formula (6). $\square$

4. ESTIMATION OF THE AVERAGE CAUSAL EFFECT

4.1. Nonparametric two-stage least squares estimator

Theorem 3 gives the nonparametric identification formulae at the population level. Based on (6), we propose a nonparametric two-stage least squares estimator of τ with finite samples $(A_i, R_i, X_{R_i}, Y_i)_{i=1}^n$. Estimation of τ_{ATT} is similar in spirit and hence omitted. We can use standard nonparametric or machine learning methods to estimate $\tau(X)$, $\text{pr}(A = a, R = 1_p)$ and

$f(X \mid A = a, R = 1_p)$; let $\hat{\tau}(X)$, $\hat{\text{pr}}(A = a, R = 1_p)$ and $\hat{f}(X \mid A = a, R = 1_p)$ denote the respective estimators. Therefore, the key is to estimate $\text{pr}(R = 1_p \mid A = a, X)$ or, equivalently, $\xi_{ra}(X)$ based on (3).

In the first stage, we obtain $\hat{f}(X_r, Y, R = r \mid A = a)$ and $\hat{f}(X, Y, R = 1_p \mid A = a)$ as the nonparametric sample analogues of $f(X_r, Y, R = r \mid A = a)$ and $f(X, Y, R = 1_p \mid A = a)$. Substituting these estimates into (3) leads to

$$\hat{f}(X_r, Y, R = r \mid A = a) = \int \xi_{ra}(X) \hat{f}(X, Y, R = 1_p \mid A = a) \, d\nu(X_{\bar{r}}), \quad (8)$$

which is a Fredholm integral equation of the first kind. Solving (8) presents several challenges. First, although Theorem 2 states that the population equation (3) has a unique solution, the sample equation (8) may not have a unique solution. Second, $\xi_{ra}(X)$ is an infinite-dimensional parameter, and its estimation often relies on some approximation. Third, solving for $\xi_{ra}(X)$ from (8) is an ill-conditioned problem, in the sense that even a slight perturbation of $\hat{f}(X_r, Y, R = r \mid A = a)$ and $\hat{f}(X, Y, R = 1_p \mid A = a)$ can lead to a large variation in the solution for $\xi_{ra}(X)$. As a result, replacing $f(X_r, Y, R = r \mid A = a)$ and $f(X, Y, R = 1_p \mid A = a)$ in (3) by their consistent estimators does not necessarily yield a consistent estimator of $\xi_{ra}(X)$ (Darolles et al., 2011).

To deal with these issues, we use a series approximation (Kress et al., 1999; Newey & Powell, 2003) in the second stage. Let the set $\mathcal{H}_J = \{h^j(X) = \exp(-X^T X) X^{\lambda_j} : j = 1, \dots, J\}$ form a Hermite polynomial basis, where $X^{\lambda_j} = X_1^{\lambda_{j1}} \cdots X_p^{\lambda_{jp}}$ with $\lambda_j = (\lambda_{j1}, \dots, \lambda_{jp})$ and $|\lambda_j| = \sum_{l=1}^p \lambda_{jl}$ increasing in j . Let $\tilde{X} = \Sigma^{-1/2}(X - \mu)$ be a standardized version of X , where μ and Σ are a constant vector and matrix. We approximate $\xi_{ra}(X)$ by $\xi_{ra}(X) \approx \sum_{j=1}^J \beta_{ra}^j h^j(\tilde{X})$. Thus, for each missing pattern $R = r$, we approximate (3) by

$$\begin{aligned} f(X_r, Y, R = r \mid A = a) &\approx \sum_{j=1}^J \beta_{ra}^j \int h^j(\tilde{X}) f(X, Y, R = 1_p \mid A = a) \, d\nu(X_{\bar{r}}) \\ &= \sum_{j=1}^J \beta_{ra}^j H_{ra}^j(X_r, Y) f(X_r, Y, R = 1_p \mid A = a), \end{aligned} \quad (9)$$

where the conditional expectation $H_{ra}^j(X_r, Y) = E\{h^j(\tilde{X}) \mid A = a, X_r, Y, R = 1_p\}$ is over the distribution $f(X_{\bar{r}} \mid A = a, X_r, Y, R = 1_p)$.

We need the empirical versions of $H_{ra}^j(X_r, Y)$ and $f(X_r, Y, R = 1_p \mid A = a)$ for estimation. First, for unit i , let $\hat{H}_{ra,i}^j = \hat{E}\{h^j(\tilde{X}) \mid A_i = a, X_{r,i}, Y_i, R_i = 1_p\}$ be a nonparametric estimator of the conditional expectation. Second, we obtain $\hat{f}(X_r, Y, R = 1_p \mid A = a)$, a nonparametric estimator of $f(X_r, Y, R = 1_p \mid A = a)$. Although we obtain these estimators based on the complete cases, we still need to partition the confounders into $(X_r, X_{\bar{r}})$ based on the missing pattern $R = r$. Because the sample version of the approximation (9) is linear, we can estimate the β_{ra}^j by minimizing the residual sum of squares

$$\sum_{i=1}^n I(R_i = r) \left\{ \hat{f}(X_{r,i}, Y_i, R_i = r \mid A_i = a) - \sum_{j=1}^J \beta_{ra}^j \hat{H}_{ra,i}^j \hat{f}(X_{r,i}, Y_i, R_i = 1_p \mid A_i = a) \right\}^2. \quad (10)$$

To ensure the estimates from (10) are well-behaved asymptotically, we need a large number of observations for each pattern $r \in \mathcal{R}$. To solve the ill-conditioned problem, we restrict the parameter space of $\xi_{ra}(X)$ to a compact space, which effectively regularizes the problem, making it well-posed. Given the approximation of $\xi_{ra}(X)$, we require the vector of coefficients β_{ra} , the concatenation of $(\beta_{ra}^1, \dots, \beta_{ra}^J)$, to satisfy $\beta_{ra}^\top \Lambda \beta_{ra} \leq B$, where Λ is a positive-definite $J \times J$ matrix and B is a positive constant. Therefore, we propose to estimate β_{ra} by minimizing (10) subject to the constraint $\beta_{ra}^\top \Lambda \beta_{ra} \leq B$. More details of the regularization are presented in the Supplementary Material.

We then estimate $\xi_{ra}(X)$ and the probability $\text{pr}(R = 1_p \mid A = a, X)$ by

$$\hat{\xi}_{ra}(X) = \sum_{j=1}^J \hat{\beta}_{ra}^j h^j(\tilde{X}), \quad \hat{\text{pr}}(R = 1_p \mid A = a, X) = \left\{ 1 + \sum_{r \neq 1_p} \hat{\xi}_{ra}(X) \right\}^{-1}$$

and finally estimate τ by

$$\hat{\tau} = \sum_{a=0}^1 \hat{\text{pr}}(A = a, R = 1_p) \int \hat{\tau}(X) \frac{\hat{f}(X \mid A = a, R = 1_p)}{\hat{\text{pr}}(R = 1_p \mid A = a, X)} d\nu(X). \quad (11)$$

We now comment on some subtle technical issues in implementing the above estimator. First, we standardize the confounders by $\tilde{X} = \Sigma^{-1}(X - \mu)$ for numerical stability. We choose μ and Σ to be the mean and covariance matrix of confounders for the complete cases. This choice is innocuous because $\mathcal{H}_J$ remains the same for other values of μ and Σ . Second, we use the importance sampling technique to approximate the integral in (11), because it is difficult to directly sample from the nonparametric density estimators. Third, we use the bootstrap to construct confidence intervals. Newey (1997) proposed a relatively simple variance estimation approach that treats the nonparametric estimators as if they were parametric given the fixed tuning parameters. For all bootstrap samples we use the same tuning parameters, such as the smoothing parameter in the smoothing splines and the bandwidth in the kernel density estimator. In the Supplementary Material, we give more technical details and illustrate the procedure with an example involving a scalar confounder.

4.2. Parametric estimation: likelihood-based and Bayesian inferences

The nonparametric estimator above suffers from the curse of dimensionality. We propose a parametric approach for moderate- or high-dimensional covariates. Let $Z_i = (A_i, X_i, Y_i, R_i)$ be the complete data and $Z_{R,i} = (A_i, R_i, X_{R,i}, Y_i)$ the observed data for unit i . The complete-data likelihood is $L(\theta \mid Z_1, \dots, Z_n) = \prod_{i=1}^n f(Z_i; \theta)$, where $\theta = (\alpha, \beta_0, \beta_1, \eta_0, \eta_1, \lambda)$ and

$$f(Z_i; \theta) = \text{pr}(R_i \mid A_i, X_i; \eta_{A_i}) f(Y_i \mid A_i, X_i; \beta_{A_i}) \text{pr}(A_i \mid X_i; \alpha) f(X_i; \lambda). \quad (12)$$

The observed-data likelihood is

$$L(\theta \mid Z_{R,1}, \dots, Z_{R,n}) = \prod_{i=1}^n \left\{ \sum_{r \in \mathcal{R}} I(R_i = r) \int f(Z_i; \theta) d\nu(X_{\bar{r},i}) \right\}.$$

Under Assumptions 6–8 as in Theorem 2, θ is identifiable if the parametric models in (12) are not overparameterized. The bounded completeness condition holds for many commonly used models,

such as generalized linear models and a location family of absolutely continuous distributions with compact support; see [Blundell et al. \(2007\)](#), [Hu & Shiu \(2018\)](#) and the Supplementary Material for additional examples. Moreover, parametric assumptions can further help to identify the model parameters even without the bounded completeness assumption. We illustrate this later.

We first discuss likelihood-based inference. Let $\tau(X_i; \theta) = E(Y_i | A_i = 1, X_i; \beta_1) - E(Y_i | A_i = 1, X_i; \beta_0)$ be the covariate-specific average causal effect, and let

$$\hat{\tau}(\theta) = n^{-1} \sum_{i=1}^n \tau(X_i; \theta), \quad \tau = \tau(\theta) = E\{\tau(X_i; \theta)\} = E\{\hat{\tau}(\theta)\}.$$

We first obtain the maximum likelihood estimate $\hat{\theta}$ and then estimate τ by $\tau(\hat{\theta})$. The formula $\tau(\theta)$ involves integrating over the distribution of the confounders. To avoid this complexity, we use $\hat{\tau}(\hat{\theta})$ to estimate τ . The bootstrap can be used to construct confidence intervals.

Next, we discuss Bayesian inference. Suppose that we can simulate the posterior distributions of the missing confounders and the parameter θ . These further induce posterior distributions of $\hat{\tau}(\theta)$ and $\tau = \tau(\theta)$. Technically, the posterior distribution of $\hat{\tau}(\theta)$ is different from that of τ . The former depends on the observed confounder values, but the latter does not. See [Ding & Li \(2018\)](#) for more discussion.

We give more computational details in the Supplementary Material, including a fractional imputation algorithm ([Yang & Kim, 2016](#)) and a Bayesian procedure for a parametric model. In future work we will develop multiple-imputation methods under Assumptions 6–8. From (12), we need to use both treatment and outcome models in the imputation step as in the full Bayesian procedure.

5. SIMULATION

5.1. Design of the simulation

We use simulation to compare our estimators with existing ones. First, we consider the unadjusted estimator, which is the simple difference-in-means of the outcomes between the treated and control groups. We use it to quantify the degree of confounding. Second, we consider the generalized propensity score weighting estimator, with the generalized propensity scores estimated separately by a logistic regression for each missing pattern ([Rosenbaum & Rubin, 1984](#)). Third, we consider three multiple-imputation estimators. The first uses the outcome in the imputation model, but the second does not ([Mitra & Reiter, 2011](#)); the third estimator uses the missingness pattern in the propensity score model ([Qu & Lipkovich, 2009](#)).

We evaluate the finite-sample performance of these estimators with the missingness of confounders satisfying Assumption 6. In the first setting, in § 5.2, one confounder has missing values and we investigate the performance of the proposed nonparametric estimator and the sensitivity to the choice of tuning parameters. In the second setting, in § 5.3, multiple confounders have missing values and we investigate the performance of the proposed parametric estimator. In each setting, we choose the sample size to be $n = 400, 800$ and 1600 , and we generate 2000 Monte Carlo samples for each sample size. For the multiple-imputation estimators, we generate 100 imputed datasets. For all estimators, we use the bootstrap with 500 replicates to estimate the variances.

5.2. One confounder subject to missingness

The confounders $X_i = (X_{1i}, X_{2i})$ follow $X_{1i} \sim N(1, 1)$ and $X_{2i} \sim \text{Ber}(0.5)$. The potential outcomes follow $Y_i(0) = 0.5 + 2X_{1i} + X_{2i} + \epsilon_i(0)$ and $Y_i(1) = 3X_{1i} + 2X_{2i} + \epsilon_i(1)$, where

Table 1. Simulation results: bias ($\times 10^{-2}$) and variance ($\times 10^{-3}$) of the point estimator of τ , variance estimate ($\times 10^{-3}$), and coverage (%) of 95% confidence intervals

Method	Bias	Var	VE	Cvg	Bias	Var	VE	Cvg	Bias	Var	VE	Cvg
(a) Comparing the nonparametric estimator with existing estimators												
	$n = 400$				$n = 800$				$n = 1600$			
Unadj	-127.5	77.4	73.7	0.3	-127.4	38.0	37.5	0.0	-127.2	17.5	18.6	0.0
GPSW	-55.1	42.4	44.2	22.2	-54.9	20.9	20.7	5.8	-54.4	9.5	9.9	0.4
MI1	41.5	35.4	36.7	40.6	41.0	15.5	17.2	9.5	40.8	7.6	8.3	0.5
MI2	-10.8	60.0	63.8	91.4	-9.2	28.8	30.8	91.4	-9.1	13.7	14.9	86.6
MIMP	29.3	73.5	71.5	83.7	28.5	33.7	32.6	65.0	28.3	14.9	16.0	30.8
NonPara	1.2	19.4	18.8	95.1	0.9	9.6	8.1	95.2	0.8	3.9	3.8	94.9
(b) Comparing the parametric estimator with existing estimators												
	$n = 400$				$n = 800$				$n = 1600$			
Unadj	32.2	85.2	85.8	81.5	32.2	44.3	42.9	65.8	31.9	20.3	21.6	43.1
GPSW	8.4	174.6	246.1	97.2	8.8	84.2	94.2	94.9	8.3	40.0	44.0	92.4
MI1	7.7	180.5	238.0	96.1	7.1	93.5	106.4	95.2	6.9	47.5	54.8	93.4
MI2	3.0	162.1	209.9	97.3	3.1	84.2	94.1	95.8	2.6	42.8	49.1	94.6
MIMP	12.9	177.0	239.2	95.7	12.2	93.9	107.5	93.8	12.1	47.4	55.0	91.8
Para	1.6	95.4	95.4	95.3	0.4	48.3	48.0	95.0	0.0	23.0	24.2	95.4

Var, variance of the point estimator of τ ; VE, variance estimate; Cvg, coverage of 95% confidence intervals; Unadj, the unadjusted estimator; GPSW, the generalized propensity score weighting estimator; NonPara, the proposed nonparametric estimator; Para, the proposed parametric estimator; for the multiple-imputation estimators, MI1 uses the outcome in the imputation, MI2 does not use the outcome in the imputation, and MIMP is the multiple-imputation missingness pattern method of [Qu & Lipkovich \(2009\)](#).

$\epsilon_i(0) \sim N(0, 1)$ and $\epsilon_i(1) \sim N(0, 1)$. The average causal effect τ is 1. The treatment indicator A_i follows $\text{Ber}(\pi_i)$, where $\text{logit}(\pi_i) = 1.25 - 0.5X_{1i} - 0.5X_{2i}$. The missing indicator of X_{1i} , R_{1i} , follows $\text{Ber}(p_i)$, where $\text{logit}(p_i) = -2 + 2X_{1i} + A_i(1.5 + X_{2i})$. The average response rate is about 67%. Other variables do not have missing values.

For the proposed nonparametric estimator, we estimate $\hat{\tau}(X)$ using cubic splines with five knots and estimate the density functions using kernel-based estimators with the Gaussian kernel. We use ten-fold crossvalidation to choose the smoothing parameters in the smoothing spline estimator and the bandwidths in the kernel-based estimators. For $\hat{\xi}_{ra}(X)$, we choose $J = 5$ Hermite polynomial basis functions and $B = 50$ as the bound for regularization.

Table 1(a) compares the nonparametric estimator with the existing estimators. The unadjusted estimator, the propensity score weighting estimator and multiple-imputation estimators are biased. As a result, the coverage rates of the confidence intervals for these methods are quite poor. Our proposed method has negligible biases and good coverages, with variances decreasing with the sample size.

To assess the sensitivity of the nonparametric estimator to the choice of the tuning parameters J and B , we specify a 4×3 design with $(J, B) \in \{(3, 50), (3, 100), (5, 50), (5, 100)\}$ and $n \in \{400, 800, 1600\}$. Table 2 shows the mean squared errors. For each (J, B) , the mean squared error decreases with the sample size. The mean squared error decreases with J , is relatively insensitive to the choice of B , and remains small across all cases.

5.3. Multiple confounders subject to missingness

Let $X_i = (X_{1i}, \dots, X_{6i})$. We generate X_{1i} and X_{2i} from $N(1, 1)$, X_{3i} and X_{4i} from $\{\text{Ber}(0.5) - 0.5\}/0.5$, $X_{5i} = X_{1i} + X_{2i} + X_{3i} + X_{4i} + \epsilon_{5i}$ with $\epsilon_{5i} \sim N(0, 1)$, and X_{6i} from $\text{Ber}(p_{6i})$ with $\text{logit}(p_{6i}) = -X_{5i}$. The potential outcomes follow $Y_i(0) = (1, X_i^T)\beta_0 + \epsilon_i(0)$ and $Y_i(1) =$

Table 2. *Simulation results for different tuning parameters: mean squared errors ($\times 10^{-3}$) of the proposed estimator of τ for different choices of (J, B) based on 2000 Monte Carlo samples*

(J, B)	$n = 400$	$n = 800$	$n = 1600$
(3, 50)	26.8	13.9	8.3
(3, 100)	27.0	14.1	8.7
(5, 50)	19.5	9.7	4.1
(5, 100)	21.3	10.2	4.5

$(1, X_i^T)\beta_1 + \epsilon_i(1)$, where $\beta_0 = (-1.5, 1, -1, 1, -1, 1, 1)^T$, $\beta_1 = (0, -1, 1, -1, 1, -1, -1)^T$, $\epsilon_i(0) \sim N(0, 1)$ and $\epsilon_i(1) \sim N(0, 1)$. The average treatment effect is $\tau = -0.5$. The treatment indicator A_i follows $\text{Ber}(\pi_i)$, where $\text{logit}(\pi_i) = (1, X_i^T)\alpha$ and $\alpha = 0.5 \times (2, 1, 1, 1, 1, -2, -2)^T$. Covariates X_{5i} and X_{6i} have missing values, but the other variables do not. The missingness pattern for X_{5i} and X_{6i} , $R_i = (R_{5i}, R_{6i}) \in \{(11), (10), (01), (00)\}$, follows a multinomial distribution with parameters $(p_{11,i}, p_{10,i}, p_{01,i}, p_{00,i})$ where

$$\text{logit}(p_{11,i}) = [1 + 3 \exp\{(1, A_i, X_i^T)\eta\}]^{-1}, \quad \text{logit}(p_{kl,i}) = [\exp\{-(1, A_i, X_i^T)\eta\} + 3]^{-1}$$

for $kl \in \{10, 01, 00\}$, with $\eta = 0.25 \times (-4, 1, 1, 1, 1, 1, -1, -1)^T$. The average percentages of these missingness patterns are about 49%, 17%, 17% and 17%, respectively.

Table 1(b) compares the parametric maximum likelihood estimator with the existing estimators. The unadjusted estimator has large biases due to confounding. The multiple-imputation estimators have large biases, although the coverages of confidence intervals seem good due to the overestimation of variances. In contrast, our estimator has negligible biases and good coverages.

6. APPLICATION

6.1. The causal effect of smoking on blood lead level

We use a dataset from the 2015–2016 U.S. National Health and Nutrition Examination Survey to estimate the causal effect of smoking on blood lead level (Hsu & Small, 2013). The dataset includes 2949 adults, consisting of 1102 smokers, denoted by $A = 1$, and 1847 nonsmokers, denoted by $A = 0$. All subjects were at least 15 years old and had no tobacco use besides cigarette smoking in the previous five days. The outcome Y is the lead level in blood, ranging from 0.05 to 23.51 $\mu\text{g/dl}$. The confounders X include the income-to-poverty level ratio, age and gender. The income-to-poverty level ratio has missing values, but the other variables do not. The missingness of income-to-poverty level is likely to be not at random because subjects with high incomes may be less likely to disclose their income information (Davern et al., 2005). It is plausible that Assumption 6 holds, i.e., that this missingness is unrelated to the blood lead level after controlling for income information. The missing rate of income-to-poverty level is 14.0% for smokers and 15.2% for nonsmokers. We apply the proposed procedure to obtain estimates separately for groups stratified by age and gender, and then average over the empirical distribution of age and gender.

Table 3(a) shows the results. Note the substantial differences in point estimates between our estimator and the competitors, illustrating the impact of the missing data assumption on causal inference in the presence of missing confounders. In contrast to the existing estimators, our estimator is better able to handle the confounders missing not at random. Based on the nonparametric estimator, smoking increases blood lead level by 0.20 $\mu\text{g/dl}$ on average.

Table 3. Results from the analysis of datasets: point estimate, standard error by the bootstrap, and 95% confidence interval

	Est	SE	95% CI		Est	SE	95% CI
(a) The causal effect of smoking on blood lead level in § 6.1							
Unadj	0.44	0.05	(0.35, 0.54)	MI1	0.34	0.05	(0.25, 0.44)
PSW	0.12	0.05	(0.02, 0.22)	MI2	0.35	0.05	(0.25, 0.44)
NonPara	0.20	0.07	(0.05, 0.36)	MIMP	0.35	0.05	(0.25, 0.44)
(b) The causal effect of education on general health satisfaction in § 6.2							
Unadj	−0.57	0.034	(−0.64, −0.51)	MI1	−0.24	0.057	(−0.36, −0.13)
GPSW	−0.25	0.054	(−0.36, −0.14)	MI2	−0.26	0.057	(−0.38, −0.15)
Para	−0.32	0.051	(−0.41, −0.21)	MIMP	−0.23	0.057	(−0.34, −0.11)

Est, point estimate; SE, standard error; CI, confidence interval; Unadj, the unadjusted estimator; GPSW, the generalized propensity score weighting estimator; NonPara, the proposed nonparametric estimator; Para, the proposed parametric estimator; for the multiple-imputation estimators, MI1 uses the outcome in the imputation, MI2 does not use the outcome in the imputation, and MIMP is the multiple-imputation missingness pattern method of [Qu & Lipkovich \(2009\)](#).

6.2. The causal effect of education on general health satisfaction

We use a dataset from the 2015–2016 U.S. National Health and Nutrition Examination Survey to estimate the average causal effect of education on general health satisfaction. The dataset includes 4845 subjects. Among them, 76% have at least high school education, denoted by $A = 1$, and 24% do not, denoted by $A = 0$. The outcome Y is the general health satisfaction score, which ranges from 1 to 5, with lower values indicating greater satisfaction. The observed outcomes have mean 2.88 and standard deviation 0.96. The confounders X include age, gender, race, marital status, income-to-poverty level ratio, and an indicator of ever having risk of prediabetes. The income-to-poverty level and prediabetes risk variables have missing values, whereas the other variables do not. The missingness of the income-to-poverty level ratio and the prediabetes risk variable is likely to be related to the missing values themselves. It is plausible that this missingness is unrelated to the outcome value conditioning on the treatment and confounders.

Table 3(b) reports the results. Although qualitatively all estimators show that education is beneficial in improving general health satisfaction, differences can be observed in the point estimates of our estimator and the competitors. This illustrates the impact of the missing data assumption on causal inference with missing confounders. Based on the parametric estimator, education improves general health satisfaction by 0.32 on average.

ACKNOWLEDGEMENT

Yang was supported in part by Oak Ridge Associated Universities, the U.S. National Science Foundation and National Institutes of Health. Wang was supported in part by the Natural Sciences and Engineering Research Council of Canada. Ding was supported in part by the U.S. National Science Foundation and the Institute of Education Sciences. The authors thank Professor Eric Tchetgen Tchetgen for valuable discussions, Professor Xiaohong Chen for useful references and the associate editor and two reviewers for helpful comments.

SUPPLEMENTARY MATERIAL

Supplementary material available at *Biometrika* online includes additional proofs, further discussions on the nonparametric and parametric estimators, and additional simulations.

REFERENCES

- AN, Y. & HU, Y. (2012). Well-posedness of measurement error models for self-reported data. *J. Economet.* **168**, 259–69.
- BLUNDELL, R., CHEN, X. & KRISTENSEN, D. (2007). Semi-nonparametric IV estimation of shape-invariant Engel curves. *Econometrica* **75**, 1613–69.
- CROWE, B. J., LIPKOVICH, I. A. & WANG, O. (2010). Comparison of several imputation methods for missing baseline data in propensity scores analysis of binary outcome. *Pharm. Statist.* **9**, 269–79.
- D'AGOSTINO JR, R. B. & RUBIN, D. B. (2000). Estimating and using propensity scores with partially missing data. *J. Am. Statist. Assoc.* **95**, 749–59.
- DAROLLES, S., FAN, Y., FLORENS, J.-P. & RENAULT, E. (2011). Nonparametric instrumental regression. *Econometrica* **79**, 1541–65.
- DAVERN, M., RODIN, H., BEEBE, T. J. & CALL, K. T. (2005). The effect of income question design in health surveys on family income, poverty and eligibility estimates. *Health Serv. Res.* **40**, 1534–52.
- D'HAUTFOEUILLE, X. (2011). On the completeness condition in nonparametric instrumental problems. *Economet. Theory* **27**, 460–71.
- DING, P. & GENG, Z. (2014). Identifiability of subgroup causal effects in randomized experiments with nonignorable missing covariates. *Statist. Med.* **33**, 1121–33.
- DING, P. & LI, F. (2018). Causal inference: A missing data perspective. *Statist. Sci.* **33**, 214–37.
- FRANGAKIS, C. E. & RUBIN, D. B. (2002). Principal stratification in causal inference. *Biometrics* **58**, 21–9.
- HANNA-ATTISHA, M., LACHANCE, J., SADLER, R. C. & CHAMPNEY SCHNEPP, A. (2016). Elevated blood lead levels in children associated with the flint drinking water crisis: A spatial analysis of risk and public health response. *Am. J. Public Health* **106**, 283–90.
- HSU, J. Y. & SMALL, D. S. (2013). Calibrating sensitivity analyses to observed covariates in observational studies. *Biometrics* **69**, 803–11.
- HU, Y. & SHIU, J.-L. (2018). Nonparametric identification using instrumental variables: Sufficient conditions for completeness. *Economet. Theory* **34**, 659–93.
- IMBENS, G. W. & RUBIN, D. B. (2015). *Causal Inference in Statistics, Social, and Biomedical Sciences*. Cambridge: Cambridge University Press.
- KRESS, R., MAZ'YA, V. & KOZLOV, V. (1999). *Linear Integral Equations*. New York: Springer, 2nd ed.
- LEHMANN, E. L. & SCHEFFÉ, H. (1950). Completeness, similar regions, and unbiased estimation: Part I. *Sankhyā* **10**, 305–40.
- LU, B. & ASHMEAD, R. (2018). Propensity score matching analysis for causal effects with MNAR covariates. *Statist. Sinica* **28**, 2005–25.
- MATTEI, A. (2009). Estimating and using propensity score in presence of missing background data: An application to assess the impact of childbearing on wellbeing. *Statist. Meth. Appl.* **18**, 257–73.
- MITRA, R. & REITER, J. P. (2011). Estimating propensity scores with missing covariate data using general location mixture models. *Statist. Med.* **30**, 627–41.
- NEWBY, W. K. (1997). Convergence rates and asymptotic normality for series estimators. *J. Economet.* **79**, 147–68.
- NEWBY, W. K. & POWELL, J. L. (2003). Instrumental variable estimation of nonparametric models. *Econometrica* **71**, 1565–78.
- NEYMAN, J. (1923). Sur les applications de la thar des probabilités aux expériences Agaricales: Essay de principe. *Statist. Sci.* **5**, 465–72. English translation of excerpts by D. Dabrowska and T. Speed.
- PEARL, J. (1995). Causal diagrams for empirical research (with Discussion). *Biometrika* **82**, 669–88.
- QU, Y. & LIPKOVICH, I. (2009). Propensity score estimation with missing values using a multiple imputation missingness pattern (MIMP) approach. *Statist. Med.* **28**, 1402–14.
- ROSENBAUM, P. R. & RUBIN, D. B. (1983). The central role of the propensity score in observational studies for causal effects. *Biometrika* **70**, 41–55.
- ROSENBAUM, P. R. & RUBIN, D. B. (1984). Reducing bias in observational studies using subclassification on the propensity score. *J. Am. Statist. Assoc.* **79**, 516–24.
- RUBIN, D. B. (1974). Estimating causal effects of treatments in randomized and nonrandomized studies. *J. Educ. Psychol.* **66**, 688–701.
- RUBIN, D. B. (1976). Inference and missing data (with Discussion). *Biometrika* **63**, 581–92.
- RUBIN, D. B. (1987). *Multiple Imputation for Nonresponse in Surveys*. New York: Wiley.
- RUBIN, D. B. (2007). The design versus the analysis of observational studies for causal effects: Parallels with the design of randomized trials. *Statist. Med.* **26**, 20–36.
- SEAMAN, S. & WHITE, I. (2014). Inverse probability weighting with missing predictors of treatment assignment or missingness. *Commun. Statist. A* **43**, 3499–515.
- U.S. DEPARTMENT OF EDUCATION (2017). *What Works Clearinghouse: Standards Handbook, Version 4.0*. Washington, DC: Institute of Education Sciences.
- YANG, S. & KIM, J. K. (2016). Fractional imputation in survey sampling: A comparative review. *Statist. Sci.* **31**, 415–32.

[Received on 28 April 2018. Editorial decision on 27 February 2019]