

Estimating Identifiable Causal Effects on Markov Equivalence Class through Double Machine Learning

Yonghan Jung
PURDUE
UNIVERSITY

Jin Tian
IOWA STATE
UNIVERSITY

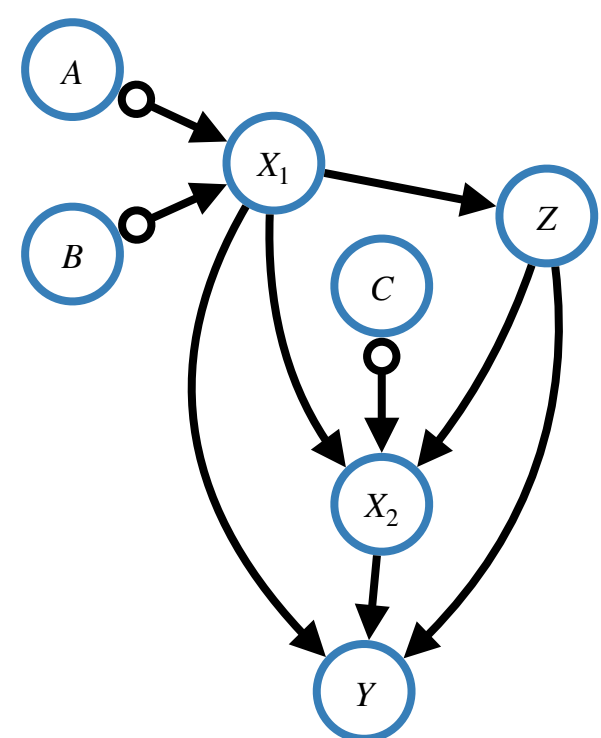
Elias Bareinboim
COLUMBIA UNIVERSITY
IN THE CITY OF NEW YORK



Summary

- ▶ Recently, *Double/Debiased Machine Learning (DML)* [1] based robust estimators have been developed for any identifiable causal effects from a *given fully specified graph* [2].
- ▶ Given observational data, we can only learn the *Markov equivalence class (MEC)* of the causal graph. A complete result for identifying causal effects from the Markov equivalence class has recently been developed [3].
- ▶ In this paper, we developed robust *DML estimators for any identifiable functional from a MEC*, which doesn't require fully specified graphs in prior.

Causal effect identification on MEC



[Figure 1] Example for Markov equivalence class

Markov equivalence class: MEC is a collection of graphs compatible with data (i.e., Two graphs G_1, G_2 are in the same MEC if the same conditional independences are implied by the graphs).

Identification on PAG: A *complete* identification algorithm (i.e., the causal effect $P(y | do(x))$ is identifiable if and only if the algorithm returns an output) has been developed [3].

[1] Chernozhukov, Victor, et al. "Double/debiased machine learning for treatment and structural parameters." *The Econometrics Journal* 21.1 (2018): C1-C68.

[2] Jung, Yonghan, Jin Tian, and Elias Bareinboim. "Estimating identifiable causal effects through double machine learning." *Proceedings of the 35th AAAI Conference on Artificial Intelligence*. 2021.

[3] Jaber, Amin, Jiji Zhang, and Elias Bareinboim. "Causal identification under Markov equivalence: Completeness results." *International Conference on Machine Learning*. PMLR, 2019.

DML for Canonical Expression

- ▶ **[Canonical Expression 1 (CE-1), Def. 3]** is given as

$$\mathcal{Q} = \sum_{\mathbf{a}} \prod_{\mathbf{B}_i \in \mathbf{C}} P(\mathbf{b}_i | \text{Pre}(\mathbf{b}_i))$$

where $\mathbf{A}, \mathbf{B}_i, \mathbf{C}$ are a subset of variables.

- ▶ An *uncentered influence function* (UIF), a key ingredient in constructing DML estimators, for CE-1 is given as

Lemma 1: UIF for CE-1

$$\mathcal{V}(\mathbf{V}; \eta = \{\theta, \omega\}) \equiv \theta_{0,1} + \sum_k \omega_k (\theta_{k,1} - \theta_{k,2})$$

where $\{\theta_{k,1}, \theta_{k,2}\}$ and $\{\omega_k\}$ are parameters represented by the conditional expectations, and can be estimated through regressions.

- ▶ The DML estimator can be constructed based on the UIF; For (D_a, D_b) are *randomly split samples* and (η_a, η_b) are nuisances trained using data (D_a, D_b) ,

$$T_N \equiv \frac{2}{N} \sum_{\mathbf{V}_{(i)} \in D_a} \mathcal{V}(\mathbf{V}_{(i)}; \eta_b) + \frac{2}{N} \sum_{\mathbf{V}_{(i)} \in D_b} \mathcal{V}(\mathbf{V}_{(i)}; \eta_a)$$

1. **Doubly Robust:** T_N converges to ψ whenever $\{\theta\}$ or $\{\omega\}$ are correct.
2. **Debiasedness:** T_N converges at $\sqrt{N}$ rate to ψ even when $\eta = \{\theta, \omega\}$ converges $N^{-1/4}$ rate.

Illustration 1: Derivation of UIF for Figure 1

A causal effect is expressed as CE-1, since

$$P(y | do(x_1, x_2)) = \sum_{z, a, b, c} P(y | x_2, x_1, z, a, b, c) P(z | x_1, a, b, c) P(a, b, c).$$

The UIF can be derived by Lemma 1 (Illustration 1 in the paper) as,

$$\mathcal{V} = \theta_{0,1} + \omega_1 (\theta_{1,1} - \theta_{1,2}) + \omega_2 (\theta_{2,1} - \theta_{2,2}),$$

where $\{\theta_{\cdot, \cdot}, \omega_{\cdot}\}$ are specified in the Illustration 1 in the paper.

DML for any identifiable causal estimands

1. **(DML-IDP in Algo. 1)** represents a causal effect as a function of CE-1 & derives/expresses a UIF for any identifiable causal effects as a function of UIFs of CE-1 in Lemma 1.

Theorem 1: Expressibility

An UIF for any identifiable causal effects can be expressed as a function of UIFs of CE-1 (in Lemma 1), through Algo. 1.

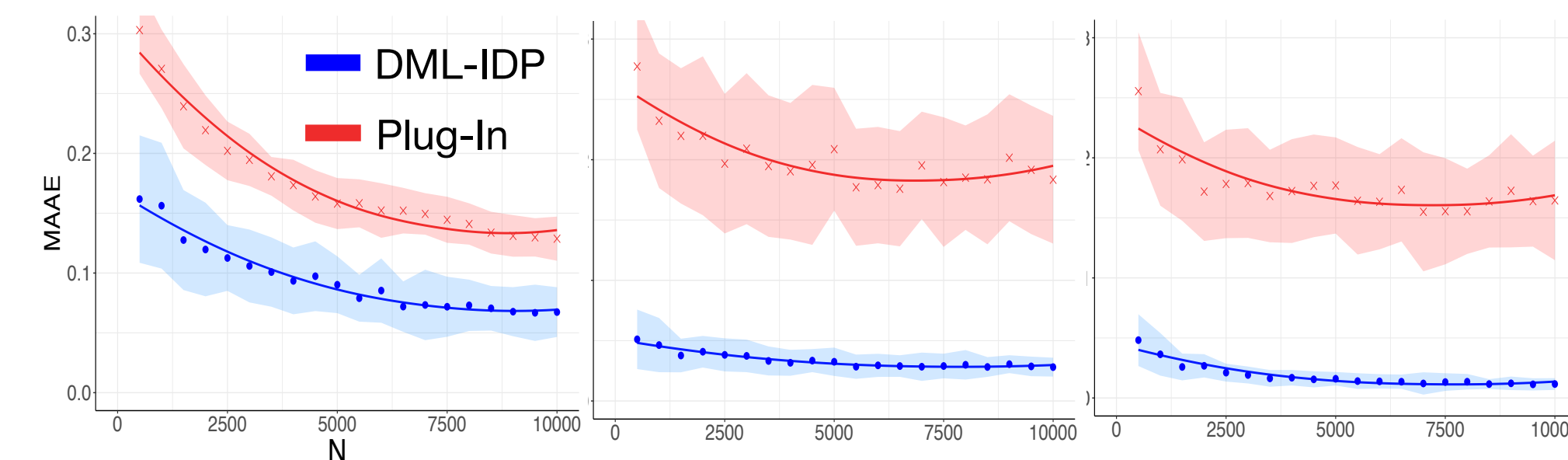
2. **(DML estimator in Def. 5)** constructs T_N based on the UIF.

Theorem 2: DML Properties

A DML estimator T_N (named DML-IDP) achieves doubly robustness and debiasedness, with respect to $\{\theta\}$ and $\{\omega\}$.

Simulation for Figure 1

- ▶ DML-IDP is compared with the plug-in estimator, the only viable estimator working for identifiable causal functional for MEC.



- ▶ **(Debiasedness; Left)** DML converges (i.e., the error 'MAAE' decreases) faster even when nuisances converge slower rate ($N^{-1/4}$).
- ▶ **(Doubly Robustness; (Center, Right))** DML converges even when models for either $\{\theta\}$ (center) or $\{\omega\}$ (right) is misspecified.

Conclusions

- ▶ **(DML-IDP)** We developed algorithms to derive corresponding UIFs for any causal effects identifiable from the MEC.
- ▶ **(DML estimator)** We introduced a general purpose causal estimators achieving *doubly robustness* and *debiasedness* properties based on the derived UIF.