



Testing policy interventions using machine learning and panel data: the Double Machine Learning method for static panel models with fixed effects

Paul Clarke and Annalivia Polselli

What we have developed

We have developed new methods for estimating the effects of treatments or policy interventions using longitudinal panel data (where repeated measures are taken of the same people over time) and machine learning (ML) algorithms.

We want to use ML because standard statistical methods for panel data require us to assume that key components of the model for the data are linear. This is generally untrue and so could result in inaccurate estimates of policy impact.

Our methods are based on so-called double/debiased machine learning (DML). This allows us to calculate valid statistical inference (e.g., standard errors and confidence intervals) for our policy effects to gauge precision as well as reduce bias.

The resulting method for panel data models with individual fixed effects is general and particularly relevant for applied researchers because panel data are widely used in applied work. The proposed procedures will broaden the reach of DML to a large family of empirical problems in various research fields for which the time dimension must be properly accounted for.

Background

A critical problem with research is that the effects of the predictor variables, required to adjust for non-random policy or treatment selection, can have nonlinear effects but the researcher does not know what form these effects take. Another problem is that omitted variables which remain fixed over time can bias estimation. We showed this method substantially outperforms traditional estimation techniques when the underlying model is highly nonlinear in the covariates, as ML algorithms successfully captured complex relationships in the data.

ML algorithms – such as the Least Absolute Shrinkage and Selection Operator (LASSO), Random Forest, Gradient Boosting, Neural Networks – are a powerful class of computational tools that allow researchers to make accurate predictions with large datasets and capture complex data structures without the need to specify the relationship between variables beforehand.

There has been an increasing interest among social scientists in using ML techniques not only for prediction purposes but also for estimating the effects of treatment or policy interventions on outputs of interest. Some authors have documented the value added of ML for causal analysis and policy evaluation in empirical works, for example, Knaus (2022), Bach et al. (2023), Strittmatter (2023), and Baiardi and Naghi (2024a,b). This is possible because researchers in the fields of econometrics, statistics and computer science have developed ML techniques for this purpose. Some researchers follow the tradition to modify ML algorithms to allow for the estimation of the causal effect (e.g., causal forests by Wager and Athey (2018)), whereas another group focuses on the use of generic ML tools to predict the models while traditional statistical techniques are used to retrieve causal effect (e.g., doubly/debiased estimators by Belloni and co-authors, 2014, 2016).

The key development, as far as our work is concerned, follows the second tradition of research and is the Double Machine Learning (DML) method developed by Chernozhukov and co-authors (2018). The DML method allows ML predictions of the functional form of the covariates (i.e., the variables that affect the outcome of interest and the treatment intervention) to be incorporated into statistical/econometric estimation. The DML method is general as it allows for the use of *any* ML algorithm (or *base learner*), and applicable to various statistical models (i.e., with exogenous, endogenous, or fully heterogeneous treatment variables). However, it was originally developed for models where the information of the unit of observation is collected in one time period only (*cross-sectional data*). Nowadays, researchers have access to richer datasets where households, firms, regions or countries are observed over repeated time periods (*panel or longitudinal data*).

The cross-sectional version of the DML method cannot be directly used with panel data because of the presence of the unobserved confounding from omitted variables, provided these variables are fixed over time (also known as the fixed effects). Ignoring the presence of these unobserved time-invariant variables, possibly correlated

with the included time-varying confounding variables, leads to incorrect estimates of the final effect. Thus, the need to provide suitable DML tools to conduct analyses with this type of data.

What did we find?

The performance of the proposed DML procedure with different machine learning algorithms is contrasted to conventional estimation technique for linear models using simulated data.

We found that our method outperforms traditional tools, especially when the relationships between variables are complicated and highly non-linear. In simpler cases, traditional methods still perform well but in practice the researcher never knows the true form of the relationship.

We also found that the choice of the machine learning algorithm the researcher uses can matter a lot, despite best-practice being followed. That is why we recommend combining several ML algorithms – a technique called *ensemble learning* – to get more reliable results.

From our results, there was often little to choose between the different approaches for handling the fixed effects. These approaches are called correlated random effects, first-differencing and within-group. This was not unexpected because the data in the simulation study were generated to satisfy conditions under which all three would perform well. But, more generally, it is first-differencing that is the most robust procedure because it imposes the fewest constraints on the distribution of the fixed effects.

Finally, we tested our method on a real-world case: the introduction of the UK's National Minimum Wage and its effect on voting behaviour. We used data from the British Household Panel Survey (BHPS). Our findings supported what we saw in our simulation exercise: our method produced more stable and trustworthy results, and showed how using complex tools in the correct way can change what we learn from data.

Uses of this research and plans for future research

Our DML method for panel data models with individual fixed effects is general and particularly relevant for applied researchers interested in estimating 'causal' effects of all kinds. Because panel data are widely used in applied work, the proposed procedures may broaden the reach of DML to a large family of empirical problems in various research fields for which the time dimension must be properly accounted for.

The DML method for panel data developed in this study can be directly applied by researchers following the instructions of the R package currently available on our GitHub repository at <https://github.com/POLSEAN/XTDML>.

Our future work will focus on extending the applicability of DML to other panel data settings that are common in many applications in the social sciences, such as, static panel data with endogenous treatment and dynamic panel data.

We will then move on to the estimation of heterogeneous treatment effects within the DML framework to estimate the intensity of the effect of a policy intervention among groups.

This has the potential to help policy-makers target interventions to those who need it the most in the most effective way.

Read the full paper

<https://academic.oup.com/ectj/advance-article/doi/10.1093/ectj/utaf011/8120202?login=false>.

Cite

Clarke, P. S. and Polsell, A. (2025). 'Double machine learning for static panel models with fixed effects', *The Econometrics Journal*, 2025; utaf011, <https://doi.org/10.1093/ectj/utaf011>

DOI

DOI: 10.1093/ectj/utaf011

References

- Bach, P., Chernozhukov, V. and Spindler, M. (2023). Heterogeneity in the US gender wage gap. *Journal of the Royal Statistical Society Series A: Statistics in Society* 187, 209-30.
- Baiardi, A. and Naghi, A. A. (2024a). 'The effect of plough agriculture on gender roles: a machine learning approach'. *Journal of Applied Econometrics* 39, 1396-402.
- Baiardi, A. and Naghi, A. A. (2024b). The value added of machine learning to causal inference: evidence from revisited studies. *Econometrics Journal* 27, 213-34.
- Belloni, A., Chernozhukov, V. and Hansen, C. (2014). 'High-dimensional methods and inference on structural and treatment effects'. *Journal of Economic Perspectives* 28, 29-50.
- Belloni, A., Chernozhukov, V., Hansen, C. and Kozbur, D. (2016). 'Inference in high-dimensional panel models with an application to gun control'. *Journal of Business and Economic Statistics* 34, 590-605.
- Chernozhukov, V., Chetverikov, D., Demirer, M., Dufo, E., Hansen, C., Newey, W. and Robins J. (2018). 'Double/debiased machine learning for treatment and structural parameters'. *Econometrics Journal* 21, C1-68.
- Knaus, M. C. (2022). 'Double machine learning-based programme evaluation under unconfoundedness'. *Econometrics Journal* 25, 602-27.
- Strittmatter, A. (2023). 'What is the value added by using causal machine learning methods in a welfare experiment evaluation?'. *Labour Economics* 84, 102412.
- University of Essex, Institute for Social and Economic Research (2018). *British Household Panel Survey: Waves 1-18, 1991-2009*. [data collection]. 8th Edition. UK Data Service. SN: 5151, <http://doi.org/10.5255/UKDA-SN-5151-2>.
- Wager, S. and S. Athey (2018). Estimation and inference of heterogeneous treatment effects using random forests. *Journal of the American Statistical Association* 113, 1228-42.