

Surfacing the Hidden Assumptions of the In-Sample Prediction Method

Nikhil Gupta and Matthias Lux

A. Introduction

In several recent antitrust class certification cases, plaintiffs have used a novel econometric method to claim empirical evidence of class-wide antitrust impact. This is a two-step method referred to by its proponents as “in-sample prediction.”¹ The first step involves the estimation of an aggregate overcharge via a regression. The second step uses the coefficients estimated by this regression to predict a but-for price for each at-issue transaction by class members. Proponents of this method claim that it can establish impact for individual observations (i.e., transactions). They identify an at-issue transaction as impacted when the actual price is more than the predicted but-for price.

While several courts have accepted this method to date,² we argue in this article that this method is not reliable.³ By carefully analyzing the way this method predicts individual impact, we show that, contrary to claims of the method’s proponents, the method does not, in general, predict the causal impact of the challenged conduct on a specific transaction.⁴

Rather, the method predicts a different mathematical object that differs from the causal impact of the challenged conduct for a specific transaction by the sum of two terms. The first of these terms is an error that may, under the right circumstances, be very small in large enough samples. The second of these terms does not vanish with larger sample sizes, and nothing can be said about its magnitude because it cannot be recovered from data.

¹ See, for example, McClave, J. and J. T. McClave. 2025. “Clarifying Common Misconceptions About the Two-Step Econometric Method for Establishing Common Impact.” *Antitrust Magazine* 39, 63—67; henceforth, McClave and McClave (2025).

² For example, some certification decisions include *In re Broiler Chicken Grower Antitrust Litigation*, No. 6:20-md-02977-RJS-CMR (E.D. Okla. May 15, 2024); *In re Pork Antitrust Litigation*, No. 0:21-md-02998-JRT-HB (D. Minn. November 14, 2021); *In re Packaged Seafood Products Antitrust Litigation*, No. 3:15-md-02670-DMS-MSB (S.D. Cal. July 30, 2019); *In re Air Cargo Shipping Services Antitrust Litigation*, No. 06-MD-1175 (JG)(VVP) (October 15, 2014); and *In re Korean Ramen Antitrust Litigation*, No. 3:13-cv-04115-WHO (N.D. Cal. January 19, 2017). The authors were not involved in any of these litigations.

³ There are several legal questions that arise due to the method’s unreliability. For example, in cases where courts granted class certification and Plaintiffs used the in-sample prediction method, was class certification unwarranted? As economists, we take no position on these issues. In our view, one of the challenges facing courts when evaluating the use of this method is the fact that the assumptions underpinning it are poorly understood even by its proponents. By highlighting these assumptions, we hope to shed light on this issue and empower fellow antitrust practitioners to correctly judge the method’s suitability and reliability in different settings.

⁴ In a technical sense, a common misunderstanding of this method is based on mistaking prediction for estimation. The method as commonly discussed does not target a fixed parameter but, as the name suggests, simply focuses on prediction. As such, studying its properties with the tools usually reserved for estimators is uninformative. For a general discussion of estimation versus prediction see, for example, Lehmann, E.L. and G. Caselli. *Theory of Point Estimation*. New York: Springer Verlag. 1998.

These two terms together can be thought of as prediction error. This prediction error never goes away: predictions inherently have irreducible error stemming from the fact that models do not capture all components influencing the outcome to be modelled.⁵

As we will explain, the in-sample prediction method only aligns with the causal impact of the challenged conduct for any specific transaction when the sum of these two terms is zero (discussed in more detail in Section C). As we explain in Section D, it is unlikely that either component is zero for all transactions. There is also no *a priori* reason to believe that their sum would equal zero for all transactions.

The presence of these two components renders the in-sample prediction method unreliable for at least two reasons. First, it means that the method does not reliably establish antitrust impact. Second, as explained in Section E, the unknown size of these two components in any given application means that the method can produce almost any fraction of false positives and false negatives: it might incorrectly identify a transaction as impacted when it was not and might incorrectly identify a transaction as unimpacted when it was.

B. Description of the in-sample prediction method

The in-sample prediction method for identifying impacted at-issue transactions involves two steps. The first step involves the estimation of an aggregate overcharge. This is typically done within a framework that compares at-issue prices to clean prices.

To make this concrete, consider a hypothetical antitrust class certification matter involving allegations of a price-fixing conspiracy. A typical overcharge regression used by plaintiffs involves comparing the prices for transactions by class members exposed to the conspiracy to transactions not exposed to the conspiracy, so-called “clean” prices. These clean prices can be from the same market prior to or after the alleged conspiracy, or they can be from a different contemporaneous market that has similar pricing dynamics.

This overcharge regression usually includes a dummy variable for the at-issue transactions during the class period as well as several control variables to account for other determinants of prices.⁶ As is the case with any regression, this regression only captures the observed elements that influence prices. Any unobserved elements that influence prices are captured in the regression's error term.^{7,8}

The second step takes the coefficients from the overcharge regression to predict a price for each transaction but for the challenged conduct. The predicted but-for price is the predicted price from the regression model with the conduct dummy set to zero.⁹

⁵ See, for example, Davidson, R. and J.G. MacKinnon. *Econometric Theory and Methods*. New York: Oxford University Press. 2004.

⁶ See equation (1) in the appendix.

⁷ In some instances, this regression may also involve interactions of control variables and the dummy variable of interest to allow for potentially different effects of the challenged conduct for different values of control variables. For clarity of exposition, we abstract away from such interactions but note that our arguments continue to be valid even if the regression specification in Step 1 allows the overcharge to vary with *observable* components. The reason, as alluded to earlier, is that prediction inherently involves error. In other words, the difference between the true value and the prediction will always involve an idiosyncratic error term from the perspective of an econometrician.

⁸ “[Error terms] are included in regression models because we are not able to specify all of the real-world factors that determine the value of [the outcome variable].” Davidson, R. and J. G. MacKinnon. *Econometric Theory and Methods*. New York: Oxford University Press. 2004, p. 2.

⁹ The predicted price of a particular transaction from a regression is equal to the estimate of the constant plus the coefficient estimates multiplied by the corresponding covariate values of that specific transactions. The predicted but-for price of an at-issue transaction is the predicted price with the covariate value of the treatment dummy (counterfactually) set to zero.

Following the logic of the in-sample prediction method, an at-issue transaction is impacted by the challenged conduct when its actual price exceeds the regression's predicted but-for price. It is then often claimed that a class member is impacted by the challenged conduct if they have at least one impacted at-issue transaction.

C. The assumptions required to recover the causal effect of the challenged conduct for an individual transaction

Reliably calculating the impact of the challenged conduct on an individual transaction requires estimating its causal effect. However, as Joshua Angrist and Guido Imbens explained in their lecture when accepting the Nobel prize:

"We cannot estimate the causal effect for an individual, without making strong assumptions."¹⁰

Using the standard framework for causal inference in economics—the potential outcomes framework—we highlight the assumptions that must hold for the in-sample prediction method to reliably predict the but-for price for a single transaction and thus recover the impact of the challenged conduct on an individual transaction.¹¹

C.1. The potential outcomes framework. The model in Step 1 of the in-sample prediction method can be thought of as a model of potential outcomes. The idea is that, for each transaction, there are two *potential* outcomes: the transaction price absent the challenged conduct and the transaction price under the challenged conduct. These outcomes are labeled as potential because the framework conceptualizes prices with and without the challenged conduct.

Note that conceptualizing the in-sample prediction method within the potential outcomes framework in Step 1 is in no way limiting. Although we believe that arguing over causal effects in regression models is easier using the potential outcomes framework, it is not necessary to use this framework to highlight the implicit assumptions of the in-sample prediction method. Other articles discussing this method sometimes choose a model that directly specifies the transaction-level impact.¹² But carefully investigating these alternative model formulations leads to the same conclusion that we present in this article.¹³

Under the assumption of a linear model for prices, there is one linear model for each of the two potential outcomes. These models describe how the potential outcomes are determined from observed determinants of prices like demand and supply factors, exposure to the challenged conduct usually measured by a single dummy variable constant across transactions and time, and an unobserved determinant of price, also called the error term.¹⁴

Impact on an individual transaction or, in other words, the causal effect of the challenged conduct on the price of a single transaction, is then simply the difference between the two potential outcomes for a given transaction: the price under the challenged conduct and the price absent the challenged conduct, the but-for price. We depict this in Exhibit 1.

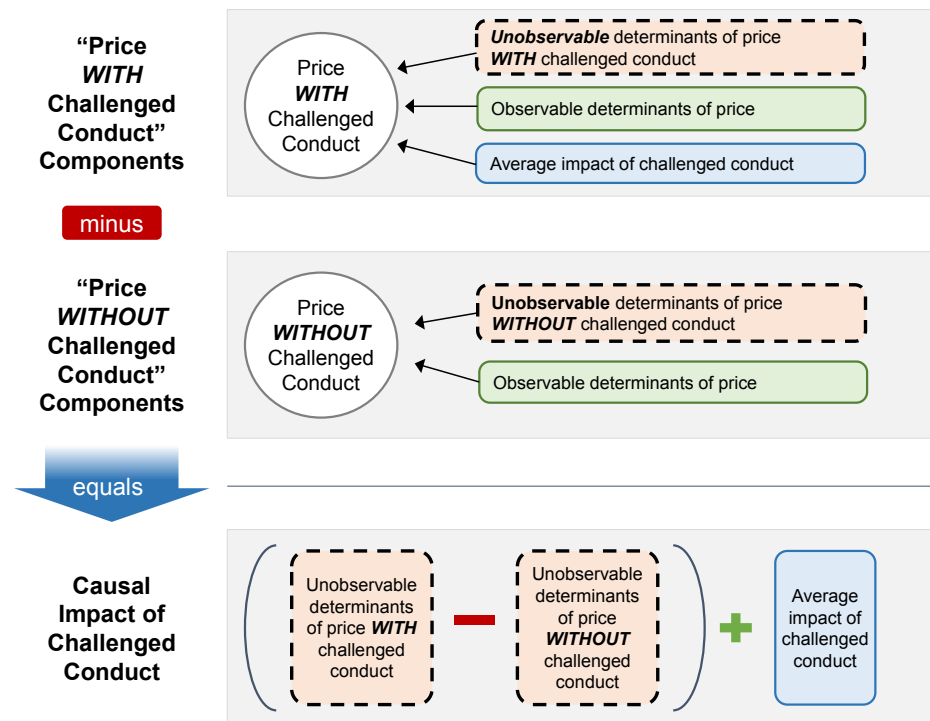
¹⁰ [FN_0725] The Committee for the Prize in Economic Sciences in Memory of Alfred Nobel, "Answering Causal Questions Using Observational Data," The Royal Swedish Academy of Sciences, October 11, 2021, p. 9; The Nobel Prize, "All Prizes in Economic Sciences," available at <https://www.nobelprize.org/prizes/lists/all-prizes-in-economic-sciences/>.

¹¹ For an application of this framework to the identification of the average causal effect of challenged conduct, see McCrary, J. and D. L. Rubinfeld. 2014. "Measuring Benchmark Damages in Antitrust Litigation." *Journal of Econometric Methods*, 3(1), 2014, 63—74.

¹² See, for example, An, Y. "Estimating Common Impact in Class Action Litigation: A Two-Step Method," *International Studies of Economics*, 2025, 20: 226—235.

¹³ See section F.3 in the appendix for a technical justification.

¹⁴ See equations (2a) and (2b) in the appendix.

Exhibit 1: The Causal Impact of the Challenged Conduct on a Transaction

As Exhibit 1 shows, given the assumption that the observable determinants of price have the same average relationship with transaction prices with and without the challenged conduct and that the challenged conduct is captured by a single dummy variable constant across transactions and time, the causal effect of the challenged conduct on a single transaction is a combination of two terms: (1) the single dummy variable capturing the average effect of the challenged conduct, and (2) the difference in unobservable determinants of prices with and without the challenged conduct.¹⁵ Note that the only reason that impact varies across individual transactions is because of unobserved components.

C.2. Taking the potential outcomes framework to data. No one observes both potential outcomes for a single transaction because a single transaction is either exposed to the challenged conduct or not. As Donald Rubin and Nobel-prize winning econometrician Guido Imbens explain, this is known as the “fundamental problem of causal inference.”¹⁶

Thus, for any given transaction, we either observe its price when exposed to the challenged conduct or its price when not exposed to the alleged misconduct. With these observed prices, an economist can derive a regression model with which to estimate the true coefficients of the *observable* components of the assumed linear model describing the potential outcomes.¹⁷

The in-sample prediction method tries to get around the fundamental problem of causal inference using *prediction*. For at-issue transactions, economists only observe the price under the challenged conduct. As a result, the in-sample prediction method proposes to predict the missing

¹⁵ See equation (3) in the appendix.

¹⁶ “The fundamental problem of causal inference is that we can observe only one of the potential outcomes for a particular subject.” Guido W. Imbens and Donald B. Rubin, *Causal Inference for Statistics, Social, and Biomedical Sciences: An Introduction*, (New York, NY: Cambridge University Press, 2015), p. i.

¹⁷ See equation (4) in the appendix.

potential outcome—the price but-for the challenged conduct—by using the coefficient estimates from the overcharge regression model. Proponents of this method claim that a reliable prediction for the causal effect of the challenged conduct on an individual transaction is the difference in the observed price for at-issue transactions and this predicted but-for price.¹⁸

C.3. The difference between what the in-sample prediction method claims to predict and what it actually predicts. Predictions, however, have irreducible uncertainty due to the error term in the model.¹⁹ The in-sample prediction method is no different. The in-sample prediction method only correctly approximates the true causal impact of the challenged conduct on a given transaction if the following two conditions hold:²⁰

1. Condition #1: the true relationship between all covariates and prices is estimated accurately and precisely by the regression model.²¹
2. Condition #2: there is no unobserved determinant of the potential price absent the challenged conduct.²² In other words, the potential price absent the challenged conduct is *fully determined* by the observed covariates included in the regression model.

Condition #1 reflects the fact that any error from the estimation of the aggregate regression model in the first step will necessarily affect any impact prediction calculated in the second step. Condition #2 is linked to the general understanding in the causal inference literature that individual treatment effects are recoverable only under extreme assumptions on how the data were generated. As we explain in section D, there is virtually no plausible scenario where both these conditions hold.

Note that these conditions are relevant both for predicting the individual quantum of impact as well as the more modest goal of determining individual impact. In the latter case, one might argue that only the sign of the difference in actual and but-for price for any individual transaction is relevant. As we show in section E, even the sign can be severely distorted using the in-sample prediction method.

D. Conditions #1 and #2 virtually never hold

Condition #1 is about whether the linear regression model in Step 1 is correctly specified, and the coefficients are estimated with precision. A regression model may be mis-specified due to, for example, omitted variable bias or measurement error.²³ In that case, the difference between the estimated coefficients and the true coefficients does not vanish as the sample grows, and condition #1 will not be met even in very large samples.

Even if the model is correctly specified, in any finite sample regression coefficients estimate the true relationship between covariates and prices with error because of sampling variability. This variability arises from the fact that economists rarely observe the entire population but only a sample drawn from it. As sample size grows, sampling variability diminishes in relevance. Thus,

¹⁸ See equation (5) in the appendix.

¹⁹ See, for example, Davidson, R. and J.G. MacKinnon. *Econometric Theory and Methods*. New York: Oxford University Press. 2004, p. 104.

²⁰ These two conditions are sufficient but not necessary. The necessary and sufficient condition is that the sum of these two potential sources of error is equal to zero. Assessing this condition, however, is inherently impossible. The two conditions we list are closer to the discussion in, for example, McClave and McClave (2025). See equation (7) in the appendix.

²¹ Captured by the first summand in equation (7) in the appendix.

²² Captured by the second summand in equation (7) in the appendix.

²³ Omitted variable bias arises when a determinant of prices is excluded from the regression but is correlated with any of the variables included in the regression. For example, in the context of transactions in a price-fixing conspiracy that involves goods differentiated with respect to quality or with respect to their attributes, common omitted variables may include measures of quality or product attributes. Measurement error arises from a discrepancy between a variable's true value and its recorded value.

if the model is correctly specified it may be reasonable to assume that Condition #1 is at least approximately met in very large samples.

Note that a commonly cited measure of regression fit, the R-squared, is not useful in assessing whether regression coefficients perfectly measure the true relationship between covariates and prices. The R-squared simply describes how close the fitted outcomes are to the observed outcomes.²⁴ By construction, it cannot tell us whether a regression is well-specified or whether there is no sampling variability that remains unaccounted for. Most importantly, it cannot tell us whether counterfactual predictions derived from the model will be close to counterfactual outcomes.

Because the prediction of individual impact in Step 2 relies on the parameter estimates from the aggregate regression in Step 1, a violation of Condition #1 necessarily impacts the predictions of individual impact.

Condition #2 is about unexplained price variation for the potential outcome when transactions are not exposed to the challenged conduct. Unexplained price variation arises in virtually every postulated relationship between prices and their determinants. It arises from determinants affecting prices that are not included in the regression. Economists typically model such unexplained price variation as an error term in regressions.²⁵

Even if Condition #1 is satisfied, the in-sample prediction method can only correctly predict a but-for price for an individual transaction when the regression model is assumed completely to capture every single determinant of pricing for every transaction not subject to the challenged conduct. Concretely, this means that the observable components included in the regression model are the only factors influencing transaction prices not impacted by the challenged conduct. No observable determinants have been omitted, and no unobservable determinants act on prices. Most economists would argue that such an assumption is, on its face, highly implausible.

E. A simulated example

In this section, we provide a small example with simulated data. We demonstrate that the histogram usually presented by proponents of the in-sample prediction method—see, for example, McClave and McClave (2025)—is, in fact, unreliable as a gauge for the true distribution of impact across transactions. Indeed, the share of transactions misclassified by the in-sample prediction method can be very large. Exactly how large that share is depends on parameters that cannot be ascertained with the data collected by the economist because of the fundamental problem of causal inference.

Our simulations are set up in a way to focus on Condition #2. We assume that the regression model is correctly specified, and all relevant determinants are observed and included in the regression model. We also simulate a very large number of transactions to ensure that any deviations between the regression coefficients and the true relationship between covariates and prices due to sampling variability is small. In other words, Condition #1 is met approximately.

Using simulations with large sample sizes also demonstrates that these two findings do not vanish with increasing sample sizes.²⁶ Hence, our findings underline the *fundamental impossibility* of the in-sample prediction method to accurately determine harm at the transaction level.

²⁴ Cameron, A.C. and P.K. Trivedi. *Microeconometrics—Methods and Applications*. New York: Cambridge University Press. 2008, p. 287

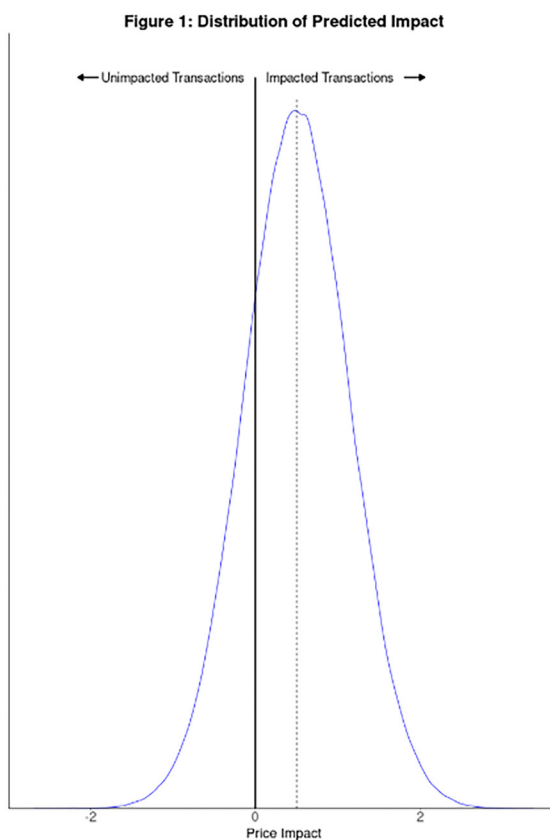
²⁵ Some econometricians consider dealing with the error term “ . . . the most important component of econometric analysis.” Wooldridge, J. *Introductory Econometrics: A Modern Approach*. South-Western College Pub. 2012, p. 4. Indeed, the idea that correct analysis of real-world data requires serious interrogation of unobserved forces and the error term is a foundational principle of econometrics. The unifying methodology of modern econometrics was articulated by Trygve Haavelmo in his seminal paper ‘The probability approach in econometrics’, *Econometrica* (1944).

²⁶ In a technical sense, these issues do not vanish asymptotically. The in-sample prediction method’s prediction is inconsistent for the true transaction-level impact.

We simulate transaction prices according to a very simple model with one observable explanatory variable influencing prices. The average impact of the challenged conduct is chosen such that prices under the challenged conduct are, on average, 10% higher than without the challenged conduct.²⁷ The parameters governing the distribution across transactions of the observable variable influencing prices are set such that all regressions have an R-squared of 0.96 or higher; high by any reasonable standard. To minimize sampling variability, we assume that there are 1,000,000 transactions, 48% of which are exposed to the challenged conduct.

We focus on simulating transactions because that clarifies exposition. In practice, it is our experience that plaintiffs claim a class member is harmed if at least one of their transactions has been impacted.²⁸ If that is the case, the in-sample prediction method is only reliable if it can correctly predict impact for an individual transaction. Moreover, if the in-sample prediction method has false positives and false negatives for calculating impact at the transaction level, it is not correct to assume that class members with many transactions are still likely to be harmed. To make that claim requires assuming that the distribution of transactions that are false positives and false negatives are not correlated across class members; there is no *a priori* reason to think this is the case.

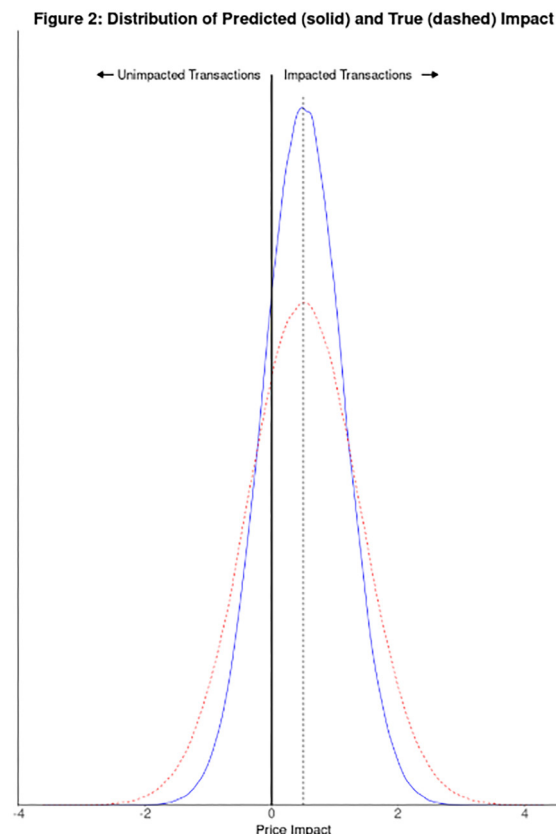
Figure 1 shows the distribution of the prediction for each transaction derived using the in-sample prediction method as a solid line. This figure is similar to the histograms presented in McClave and McClave (2025). As discussed in that article, the predictions center around the average impact, here 0.5. Some transactions have predictions below this average, and some have predictions above this average. The method predicts that a transaction is impacted if it lies to the right of the black, solid vertical line.



²⁷ See section F.4 in the appendix for the exact parameters chosen for the simulations in this section.

²⁸ See, for example, *Air Cargo* (n. 2) at 37 (“... it is enough if they provide sufficient evidence to demonstrate that substantially all class members were overcharged at least once.”).

Having simulated the data, we also know the true individual impact on each transaction, and we can compare the distribution of the in-sample prediction method's predictions to the distribution of the true transaction-specific impact. That is shown in Figure 2. We see that the distribution of the true impact (the dashed curve) is much wider than the distribution of the in-sample prediction method's predictions (the solid curve). The fact that the distribution is wider is an artifact of our simulation. With different parameters, it would be possible to have the relationship inverted. Thus, the type of figure presented in McClave and McClave (2025) does not show the true distribution of impact across transactions.



Matters are worse, in fact. One might conclude from Figure 2 that the distributions are not the same but do look quite similar. Might this mean that the in-sample prediction provides at least a good *approximation* of the true impact for a transaction?

The answer is no. Even if the distributions were even more similar than those in Figure 2, this would not tell us much about where an individual transaction lies inside this distribution. Put differently, an individual transaction could lie *at different points* on the two distributions. Hence, the in-sample prediction method could be predicting impact in a case where the true impact was zero or negative (a false positive), or the in-sample prediction method might not predict impact where the true impact was positive (a false negative). This is made clear with the results shown in Table 1.

Table 1 shows the fraction of transactions for both types of errors as a function of the relevant underlying regression parameters. Here, we use the same regression specification as for Figure 1. The only difference is that we are varying the parameters determining the unobserved components for the potential prices under the challenged conduct and without the challenged conduct. These parameters are the variances of the unobserved components of prices with and without the challenged conduct and the correlation coefficient of these two components.

Table 1: Shares of False Positives and False Negatives

Variance of the Unobservable Component...		Value of the Correlation Coefficient					
... under the challenged conduct	... but-for the challenged conduct	-0.6		0		0.6	
		False positive	False negative	False positive	False negative	False positive	False negative
1.2	1.2	11%	4%	14%	9%	14%	16%
0.4	1.2	18%	2%	21%	8%	23%	16%
1.2	0.4	8%	3%	9%	7%	7%	11%
0.4	0.4	13%	2%	14%	7%	11%	13%

The table shows that the share of transactions with false positives ranges in our simulation from about 7% to as high as 23%. The share of false negatives ranges from 2% to as high as 16%. Generally, the share of false negatives is higher when the correlation between the unobserved components determining potential prices is positive. For false positives, the trend is less clear.

From an economic point of view, the case of positive correlation is arguably the more intuitive scenario. The unobserved components determining potential prices capture elements related as well as unrelated to the challenged conduct. The shared component unrelated to the challenged conduct creates joint co-movement in the unobserved components of the two potential prices, and this joint co-movement would result in a positive correlation between the two.

In conclusion, we have made transparent the assumptions required for the in-sample prediction method to correctly measure the impact of the challenged conduct on a given transaction, and we have argued that these assumptions rarely hold. Our simulations show that when these assumptions do not hold, the method yields a potentially large number of false positives. Hence, taken together our results demonstrate the *fundamental impossibility* of the in-sample prediction method to measure the impact of the challenged conduct on a given transaction. Correspondingly, this method is unreliable and cannot show class-wide impact.

F. Mathematical Appendix

This appendix outlines the key mathematical derivations that underlie our arguments in the text.

F.1. The potential outcomes framework and causal impact. We assume that prices or outcomes are generated by

$$y_{it} = x'_{it}\beta + \delta + \varepsilon_{it} \quad (1)$$

where y_{it} is the outcome or price for transaction i at time t , x_{it} is a vector of observable components influencing outcomes—and can include a constant and other fixed effects—, δ is the impact of the challenged conduct on prices, and ε_{it} is a mean-zero residual uncorrelated with the other right-hand side components.

The potential outcomes framework typically uses the notation $y_{it}(1)$ and $y_{it}(0)$ to denote the potential outcome with and without the challenged conduct, respectively. Using this notation, the above model can be thought of as a model of potential outcomes, namely:

$$y_{it}(1) = x'_{it}\beta + \delta + u_{it} \quad (2a)$$

$$y_{it}(0) = x'_{it}\beta + v_{it} \quad (2b)$$

The causal impact of the challenged conduct on any given transaction is the difference between the potential outcome for the transaction under the challenged conduct and without the challenged conduct. Mathematically, this is written as:

$$y_{it}(1) - y_{it}(0) = \delta + u_{it} - v_{it} \quad (3)$$

Equation (3) mathematically expresses the object that the in-sample prediction method claims to reliably predict for each transaction: the causal impact of the challenged conduct on each transaction. As discussed in the main text, this causal impact depends on both the average effect of the challenged conduct δ and the two unobserved determinants of prices, u_{it} and v_{it} .

F.2. What the in-sample prediction method predicts. Because of the fundamental problem of causal inference discussed in the text, economists cannot observe both potential outcomes for any transaction. Instead, they estimate a regression model that relates to these potential outcomes. Let D_{it} be a dummy variable that takes the value 1 if the transaction is affected by the challenged conduct and zero otherwise. Then, using the potential outcomes framework, the regression equation can be written as:

$$y_{it} = D_{it} y_{it}(1) + (1 - D_{it}) y_{it}(0) = x'_{it} \beta + \delta D_{it} + \varepsilon_{it} \quad (4)$$

where y_{it} is the observed price of a transaction. The error term in this regression, ε_{it} , is similarly a composite of the true error terms, namely

$$\varepsilon_{it} = D_{it} u_{it} + (1 - D_{it}) v_{it}.$$

See McCrary and Rubinfield (2014) for a discussion under which conditions on ε_{it} and D_{it} the coefficients are identified.

Note that this is exactly equation (1), often taken as the starting point for empirical analyses of the impact of challenged conduct. Doing so, however, ignores that the error term ε_{it} is not a primitive but, rather, a composite of two error terms. And the interpretation of these two error terms highlights the assumptions made when viewing ε_{it} as a primitive. This regression equation is Step 1 of the in-sample prediction method.

The in-sample prediction method then uses this regression equation in Step 2 of the method to predict the but-for price for at-issue transactions via

$$\hat{y}_{it}(0) = x'_{it} \hat{\beta} \quad (5)$$

Using this predicted but-for price, impact on an individual transaction is predicted by taking the difference between the actual price and the model's predicted but-for price for an at-issue transaction:

$$y_{it}(1) - \hat{y}_{it}(0) = \delta + x'_{it} (\beta - \hat{\beta}) + u_{it} \quad (6)$$

Hence, the difference between the true impact from equation (3) and the prediction based on the in-sample prediction method is

$$y_{it}(1) - y_{it}(0) - (y_{it}(1) - \hat{y}_{it}(0)) = x'_{it} (\beta - \hat{\beta}) + v_{it} \quad (7)$$

F.3. Our results do not rely on using the potential outcomes framework. Using the definition of the residual ε_{it} , it is possible to rewrite (4) as

$$y_{it} = x'_{it} \beta + (\delta + u_{it} - v_{it}) D_{it} + v_{it}$$

This is essentially equal to equation (4) in An (2025) with $\Delta \beta_{it} = u_{it} - v_{it}$.²⁹ In addition, we have that $E[u_{it} - v_{it} | x_{it}, D_{it}] = 0$ by assumption, and also $E[\Delta \beta_{it} | \alpha_p, x_{it}, D_{it}] = 0$; see An (2025, p. 229) for this last equality. This demonstrates the fact that the two formulations are identical and only differ in their

²⁹ An, Y. "Estimating Common Impact in Class Action Litigation: A Two-Step Method," *International Studies of Economics*, 2025, 20: 226–235.

assumptions on the primitive: in our specification, we focus on potential outcomes as primitives, whereas An (2025) models the individual component of the treatment effect directly as $\Delta\beta_{it}$.

To see that this specification relies on a comparable assumption as Condition #2, note that the second line of equation (7) in An (2025) states that (using their notation)

$$\hat{\beta}_{it} = (\alpha_i - \hat{\alpha}_i) + X'_{it}(\gamma - \hat{\gamma}) + \beta_{it} + \varepsilon_{it}$$

Here, $\hat{\beta}_{it}$ is the same as $y_{it}(1) - \hat{y}_{it}(0)$ in our notation. It is now easy to see that for $\hat{\beta}_{it}$ to be a proper estimate of the transaction-level damage β_{it} , both the coefficients must be estimated without error (Condition #1) and ε_{it} has to be zero. Looking at equation (1) in An (2025), this means that prices both exposed to the challenged conduct and not exposed to the challenged conduct must be modelled without error, thus implying our (weaker) Condition #2.

F.4. Technical details of the simulation. We simulate prices according to (2a) and (2b) above. The covariates include a constant with value 0.5 and a single covariate varying across transactions drawn from a normal distribution with mean 1 and standard deviation 1.25. The coefficient $\beta = 4.5$ and the average causal impact of the challenged conduct is $\delta = 0.5$. The fraction of transactions exposed to the challenged conduct is 0.48.

The error terms influencing the counterfactual prices are drawn from a multivariate normal distribution according to

$$\begin{pmatrix} v_{it} \\ u_{it} \end{pmatrix} \sim N(0, \Sigma)$$

with

$$\Sigma = \begin{pmatrix} \sigma_v & \sigma_{uv} \\ \sigma_{uv} & \sigma_u \end{pmatrix}$$

where $\sigma_{uv} = \rho\sigma_v\sigma_u$. For Figure 1 and 2, we choose $\sigma_v = 1.2$, $\sigma_u = 0.4$, and $\rho = 0.6$. The values for ρ , σ_u , and σ_v for the simulations in Table 1 are as described in that table. ●