

Why This Popular Class Cert. Approach Doesn't Measure Up

By **Celeste Saravia and Daniel Ramsey** (October 24, 2025)

A key question in class certification decisions in the U.S. is whether classwide evidence can be used to determine if the challenged conduct resulted in harm to all, or almost all, proposed class members.[1]

In this article, we discuss and evaluate an approach that has been used to address this question in more than 10 recent antitrust cases.[2]

The approach, which plaintiffs experts sometimes call the in-sample prediction approach, claims to estimate the impact of the challenged conduct on each transaction.[3] Specifically, in a price-fixing case, it purports to provide an estimate of the overcharge on each transaction.[4]

These overcharge estimates could be, for example, \$3.20 for the first transaction, \$1.10 for the second transaction and so on, for every transaction.

The expert would then use these estimates of the transaction-specific overcharges to determine the percentage of class members harmed on at least one transaction.

In every class action we identified in which a plaintiffs expert used this approach, the expert found that almost all class members were harmed, and the judge certified the class, as shown in Table 1.

In light of recent usage of this approach, and its success in achieving class certification, plaintiffs experts are likely to use the approach in future cases as well.

We explain that this approach is unreliable, because it fails two fundamental tests of reliable econometric methods.

First, it is not a consistent estimator, meaning that its estimate will not converge to the true value of the overcharge as the data grows large. For example, if the true overcharge on a transaction was \$1, the approach's estimate could be quite far from \$1, even with a near-infinite number of observations.

Second, it has a false positive error rate far beyond typical thresholds, meaning that it will find that a substantial percentage of class members were harmed even if they were not harmed.

Additionally, the approach is contradicted by a large academic literature.

Table 1: Cases in Which a Plaintiffs Expert Used the In-Sample Prediction Approach



Celeste Saravia



Daniel Ramsey

Case	Class certified	Class cert decision date	Class members found harmed, per plaintiffs expert
In re: Broiler Chicken Grower Antitrust Litigation (U.S. District Court for the Eastern District of Oklahoma)	Yes	2024	95%-100%
Le v. Zuffa LLC (U.S. District Court for the District of Nevada)	Yes	2023	96.3%-98.8%
In re: Pork Antitrust Litigation (U.S. District Court for the District of Minnesota)	Yes	2023	99%
In re: Broiler Chicken Antitrust Litigation (U.S. District Court for the Northern District of Illinois)	Yes	2022	93%
In re: Packaged Seafood Products Antitrust Litigation (U.S. District Court for the Southern District of California)	Yes	2019	94.5%
In re: Capacitors Antitrust Litigation (U.S. District Court for the Northern District of California)	Yes	2018	98%
In re: Domestic Drywall Antitrust Litigation (U.S. District Court for the Eastern District of Pennsylvania)	Yes	2017	98%
In re: Korean Ramen Antitrust Litigation (U.S. District Court for the Northern District of California)	Yes	2017	98%
In re: Air Cargo Shipping Services Antitrust Litigation (U.S. District Court for the Eastern District of New York)	Yes	2014	96%
In re: Chocolate Confectionary Antitrust Litigation (U.S. District Court for the Middle District of Pennsylvania)	Yes	2012	98%

Overview of the In-Sample Approach and Illustrative Example

The in-sample approach purports to offer a two-step process for assessing which proposed class members were harmed — i.e., suffered impact.[5] In the first step, a plaintiffs expert seeks to determine the average classwide impact, if any, of the challenged conduct.

For example, in a price-fixing case, the expert seeks to estimate the average overcharge

paid by class members. The expert might purport to find that, for example, the challenged conduct caused the prices paid by class members to be \$2.00 higher on average.

The expert typically performs this first step by using transaction data to estimate a regression equation. The dependent variable is a measure of price, and the explanatory variables include multiple control variables that can explain variation in prices across transactions and an indicator variable that indicates whether the transaction occurred during the alleged conspiracy.[6]

The plaintiffs expert's claim that this regression provides an estimate of the average overcharge is only valid if several assumptions hold. However, even if the expert finds that the conduct had an impact on average prices, this does not imply that all class members were harmed, nor does it show which class members were harmed.

In the second step, the plaintiffs expert uses the first-step regression results to calculate which transactions and which class members were harmed. The following equation provides the purported transaction-specific overcharges:[7]

Transaction-specific overcharge = average overcharge + residual

For example, if the first-step regression estimates an average overcharge of \$2 across all transactions, and the first transaction has a residual of \$1 from that regression, then the technique would find that the overcharge on the first transaction is \$3.

The residual, which we discuss in more detail below, is the portion of the outcome that is not explained by the model. For example, if a transaction's price was \$12, but the first-step regression model predicted it to be \$11, then the residual is \$1.

A plaintiffs expert using this approach typically does not require the transaction-specific overcharge to meet any statistical threshold. The plaintiffs expert and court typically conclude that a class member is harmed if the class member was harmed on at least one transaction, rather than requiring the class member to be harmed on net across all of their transactions.[8]

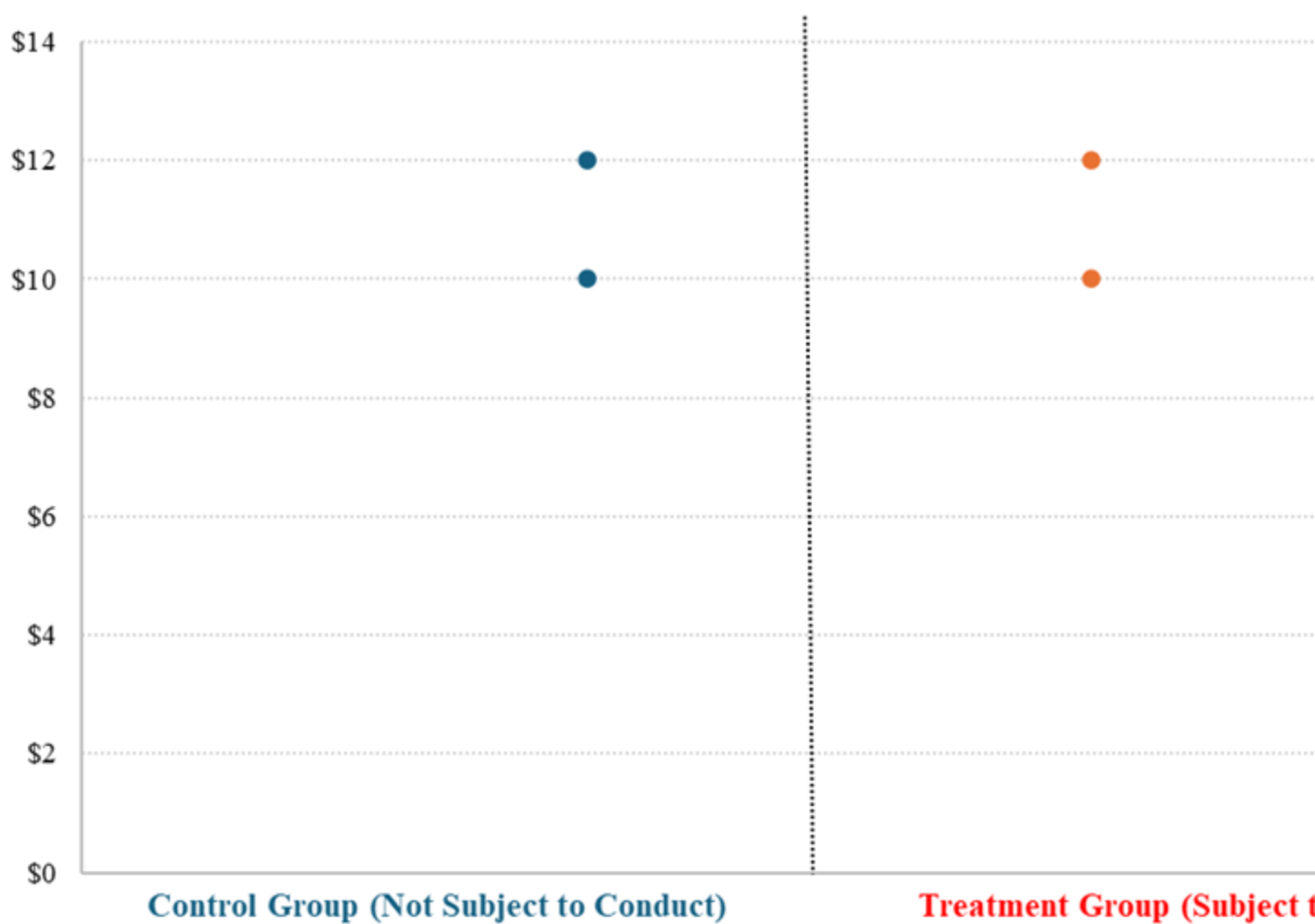
We illustrate the approach with a simple example in which there are four transactions. The data is constructed such that the conduct has no impact on the average price or on the price of any transaction.[9]

Figure 1 depicts these four transactions. Two of the transactions are subject to the challenged conduct — the treatment group, depicted in red — and these have prices of \$12 and \$10.

The other two transactions are not subject to the challenged conduct — the control group, depicted in blue — perhaps because they occurred in a different region or year, and these also have prices of \$12 and \$10. There are no covariates.

Figure 1: Example Transaction Prices

Transaction Prices for Control Group (Blue) and Treatment Group (Red)



Note: This figure depicts the four transactions in the example. The in-sample approach would find the \$12 transaction in red to have a \$1 overcharge and the \$10 transaction in red to have a \$1 undercharge.

The in-sample approach applied to this data would correctly determine that the average impact was zero, but it would incorrectly determine that one of the two transactions subject to the challenged conduct was harmed.[10]

In particular, the approach finds that the \$12 transaction would have been \$11 but for the conduct, meaning it had an overcharge of \$1. And the approach finds that the \$10 transaction benefited from the challenged conduct by \$1.

Thus, even though the prices are identical in the control group and treatment group, the approach somehow finds that both treatment group transactions were affected by the conduct — one harmed and one benefited.[11]

It obtains these results because the residual is \$1 for the \$12 transaction, and the residual is -\$1 for the \$10 transaction, which it adds to the average impact of zero.

This example illustrates two fundamental flaws of the in-sample approach. First, it does not correctly identify the overcharge for any transaction. Here, the true impact was zero for each transaction, but the approach found an overcharge of \$1 on one transaction and -\$1 on the other transaction.

Second, the approach has a substantial false positive rate, meaning that it finds a high percentage of class members harmed even if they were not harmed — 50%, in this example.

In the next section, we explain that these flaws are not limited to this example, but are fundamental properties of the approach.

Evaluation of the Reliability of the In-Sample Approach

Evaluation of Whether the Approach Is a Consistent Estimator

First, we assess whether the approach meets a minimal requirement of a reliable econometric approach: whether it is a consistent estimator of the transaction-specific overcharges.[12]

This involves assessing whether the estimate provided by the approach will converge to the true value as the size of the data increases. For example, if the true overcharge on a given transaction is \$1.00, will the approach's estimate converge to \$1.00?

If it is not a consistent estimator, the approach could provide an estimate that is substantially different from \$1.00 even as the size of the data increases to infinity, and thus, the approach would not reliably determine the harm incurred by individual transactions or class members.

We find that the approach is not a consistent estimator. As shown in the equation above, the approach obtains the transaction-specific overcharge by adding the residual to the average overcharge. The approach provides no basis for assuming the residual is entirely due to the conduct, and such an assumption is contradicted by econometric principles.

Regression residuals are the portion of the outcome — the price, in this example — that is unexplained by the model, and are determined by numerous factors, such as omitted variables or data error. Essentially every econometric model will have regression residuals, since it is typically not possible to perfectly predict every transaction's price.[13]

Since the approach includes the residual in its estimate of the transaction-specific overcharge, the overcharge estimate errantly includes factors that are unrelated to the overcharge, such as omitted variables or data error.

For example, assume that in our example above that there were an omitted variable that explains some of the variation in prices. The in-sample approach would falsely attribute the impact of that omitted variable to the conduct.

This error will persist no matter how many observations are used in the analysis. Thus, the approach is not a consistent estimator.[14]

Evaluation of the False Positive Rate of the Approach

Next, we assess whether the in-sample approach has a false positive rate within the typical thresholds. False positives are instances when the true impact of the conduct is zero, but the approach finds an impact.[15]

Common econometric and statistical approaches typically have false positive error rates of 5% or less.[16] If an approach has a high false positive rate, it is typically not a reliable approach.[17]

We find that the in-sample approach has a false positive rate well above 5%, and it can approach 100%. The equation above shows why the in-sample approach has a high false positive rate.

Consider an example in which the challenged conduct had no impact on any transaction, and thus, the average impact is also zero. Then as the number of transactions increases, the estimated average impact will approach zero, and the approach's estimated impact on a given transaction is just the regression residual.

Regression residuals are mean zero for any regression — i.e., there are both positive and negative residuals that sum to zero.[18] Thus, if the distribution of residuals is fairly symmetric, then roughly half of the residuals will be positive and roughly half will be negative.

Since, as discussed above, the approach typically does not require the estimated impacts to meet any threshold of statistical significance, the approach would therefore incorrectly conclude that roughly half of transactions are harmed — the positive residuals in a price-fixing case — even if the conduct had no impact on any transaction.

This false positive issue is exacerbated by the fact that each class member typically has many transactions, sometimes even hundreds of transactions.[19] The probability that the approach finds that a class member is harmed on at least one transaction increases with the number of transactions of that class member.

If the true impact of the conduct is zero, a class member has two transactions, and the residuals are symmetric and uncorrelated across the transactions, then the probability that the approach finds that the class member is harmed on at least one transaction is 75% — calculated as $1 - (1 - 0.5)^2$.

If the true impact of the conduct is zero, a class member has five transactions, and the residuals are symmetric and uncorrelated, then the probability that the approach finds that the class member is harmed on at least one transaction is 97% — calculated as $1 - (1 - 0.5)^5$.

Thus, even though the true impact of the challenged conduct is zero in this example, the approach would find that almost all class members were harmed.[20]

Evaluation of Whether the Approach Is Consistent With the Academic Literature

There is a body of large academic literature on causal inference, which discusses various approaches to estimating the impact of policies or events, such as alleged collusion.

This literature concludes that it is not generally possible to do what the in-sample approach claims to do, which is to estimate individual causal effects such as transaction-specific overcharges.[21]

For that reason, the academic literature typically estimates average effects, which in the case of a price-fixing matter would be the average overcharge, rather than transaction-specific overcharges.

Moreover, while there is recent academic literature that attempts to go beyond estimating average effects, it does not propose a method like the in-sample approach.[22]

Conclusion

Plaintiffs experts have used the in-sample approach in more than 10 recent antitrust class actions to argue that all or almost all proposed class members were harmed by the challenged conduct.

In each case in which the court has issued a class certification decision, the court certified the class, and at least in part, credited the plaintiffs expert's in-sample approach. However, we find that it is not a reliable approach. It fails two fundamental tests of reliable econometric methods, and is contradicted by the academic literature.

Celeste Saravia, Ph.D., is vice president and co-head of the antitrust and competition practice at Cornerstone Research.

Daniel Ramsey, Ph.D., is a principal at the firm.

Disclosure: Celeste Saravia was a testifying expert, and Daniel Ramsey was a consultant, on behalf of the defense in In re: Broiler Chicken Grower Antitrust Litigation.

The opinions expressed are those of the author(s) and do not necessarily reflect the views of their employer, its clients, or Portfolio Media Inc., or any of its or their respective affiliates. This article is for general information purposes and is not intended to be and should not be taken as legal advice.

[1] Different U.S. circuit courts have taken different positions on the meaning of "all or almost all" class members. For example, some circuits require there to be less than a de minimis number of unharmed class members, and others allow more than a de minimis number. On June 5, the U.S. Supreme Court dismissed a writ of certiorari related to this question as improvidently granted. *Lab'y Corp. of Am. Holdings v. Davis*, No. 24-304, (U.S. June 5, 2025) (Kavanaugh, J., dissenting).

[2] The points in this article are discussed in more detail in our more technical article. See Daniel Ramsey and Celeste Saravia, *Evaluating the In-Sample Prediction Approach to Assessing Class-Wide Impact for Class Certification*, (2025) (Ramsey and Saravia (2025)), available at <https://www.cornerstone.com/wp-content/uploads/2025/10/EVALUATINGTHE-IN-SAMPLE-PREDICTIONAPPROACH-TO-ASSESSING-CLASS-WIDE-IMPACT-FOR-CLASS-CERTIFICATION.pdf>.

[3] Presumably, plaintiffs experts call it the in-sample prediction approach since the experts

use the same data sample for estimating the coefficients and making the predictions.

[4] The in-sample prediction approach has been used to analyze different outcomes in different cases, such as prices in a price-fixing case or wages in a wage-fixing case. We focus on the price-fixing example, with harm measured by the price overcharge, but the same critiques would apply to the use of the in-sample method in a wage-fixing case. While our article focuses on antitrust class actions, it is possible that the in-sample prediction approach has been used in other types of class actions.

[5] This section is based on publicly available expert reports that use the in-sample approach. See Ramsey and Saravia (2025) for further detail.

[6] The plaintiffs expert typically estimates an equation like the following:

$$P_{it} = \beta X_{it} + \delta D_{it} + \varepsilon_{it}$$

P_{it} is a measure of the price paid for transaction i at date t , X_{it} is a vector of control variables for transaction i at date t , and β is the coefficient vector for the control variables (including a regression constant). D_{it} is a conduct indicator variable that equals 1 if transaction i at date t was allegedly subject to the challenged conduct and 0 otherwise. ε_{it} is the error term.

[7] In Ramsey and Saravia (2025), we show that this is an alternative mathematical expression of the in-sample approach. When plaintiffs experts explain the approach, they typically state that after estimating the first-step regression, they switch the conduct variable from $D_{it} = 1$ to $D_{it} = 0$, and then form regression predictions, which they call the but-for prediction. They then calculate the purported transaction-specific overcharge as the difference between the price and the but-for predicted price.

[8] While it is not the focus of this article, the economic basis for concluding that a class member is harmed if the class member was harmed on at least one transaction, rather than being harmed on net, is unclear.

[9] See Ramsey and Saravia (2025) for an example in which the conduct had an average impact on prices.

[10] Perhaps a plaintiffs expert would contend that they would not use the in-sample approach if they find a zero impact in the first step. However, the purpose of this example, as well as the false positive tests below, is not to determine the percent of class members who are truly harmed, but rather, to assess the reliability of the approach. See Ramsey and Saravia (2025).

[11] In practice, in determining whether a class member was harmed, the approach essentially discards transactions that were found to have benefited from the conduct, since the approach does not determine the net harm.

[12] Jeffrey M. Wooldridge, *Introductory Econometrics: A Modern Approach*, 168–169, 779 (5th ed. 2012) (Wooldridge) at 169 ("[V]irtually all economists agree that consistency is a minimal requirement for an estimator").

[13] William H. Greene, *Econometric Analysis, International Edition* (7th ed. 2012) at 53.

[14] The in-sample approach only identifies the transaction-specific overcharge if the entire residual is due to the heterogeneous impact of the conduct — i.e., how much the

overcharge for each transaction varies from the average overcharge. This would require the regression to perfectly predict the data if the impact were not heterogenous or if the transaction was not subject to the conduct, which is not plausible. See Ramsey and Saravia (2025).

[15] Wooldridge at 779.

[16] Wooldridge at 779.

[17] Wooldridge at 771; Marianne Bertrand et al., How Much Should We Trust Differences-in-Differences Estimates?, 119 Q.J. Econ. 249 (2004).

[18] Wooldridge at 36.

[19] See Ramsey and Saravia (2025).

[20] See Ramsey and Saravia (2025) for empirical simulations that further demonstrate the false positive issue.

[21] Miguel Hernan and James Robin, Causal Inference: What If 16 (2024) ("Our discussion of randomized experiments refers to population or average causal effects because individual causal effects cannot generally be identified"); Mingzhang Yin et al., Conformal Sensitivity Analysis for Individual Treatment Effects, 119 J. Am. Stat. Assoc. 122, (2024) at 3 ("Therefore, the [individual treatment effect] ITE, which requires knowing both the potential outcomes, can never be observed. Furthermore, unlike population-level causal estimands, an ITE is inherently random. Even with a known joint distribution ... an ITE is not point-identifiable").

[22] See Ramsey and Saravia (2025) for further discussion.