

Online Appendix for No Call, No Show: Impact of No-Shows on Customer Attrition in Online-to-Offline Services

Nil Karacaoglu and Simin Li

Appendix A: Additional Analyses for Sections 3 and 4

A.1. Summary Statistics

In Section 3.1, we introduce our dataset.¹ The summary statistics for our sample are as follows:

Table A.1 Summary Statistics

Variable	Mean	SD	Min	Max
R	0.33	0.47	0	1
NoShow	0.04	0.18	0	1
CustomerOrderHistory	6.41	14.69	0	247
Price	119.45	71.30	19	800
DiscountAmount	27.02	41.85	0	452

A.2. Robustness Analyses

A.2.1. Matching. Our main identification strategy leverages quasi-experimental variation in no-shows to assess their impact on customer attrition. In an ideal experiment, the platform would randomly assign no-shows to cleanings, ensuring that cleaners assigned to specific cleanings deliberately do not show up. However, such randomization is impractical because it would deliberately degrade the customer experience and the long-term success of the platform.

As demonstrated in Section 4.2.1, accurately predicting no-shows is challenging, with our XGBoost model achieving only a limited predictive performance, despite including a broad set of features that capture customer characteristics as well as other features that could impact customer attrition. This finding strengthens our identification strategy and reduces endogeneity concerns: it indicates that factors such as customer characteristics, which impact customer attrition, are not simultaneously strongly associated with the likelihood of no-shows.

To ensure further comparability between treatment (cleanings that experienced no-shows) and control cleanings (those without no-shows), we employ a matching algorithm as a robustness check. Specifically, we match cleanings that experienced no-shows to those that did not based on observable characteristics that might influence customer attrition and no-show probabilities. By minimizing covariate imbalance, our approach reduces bias from observable differences while using these variables as proxies for potential unobserved factors. This strengthens our ability to attribute differences in customer attrition to no-shows rather than to pre-existing disparities across treated and control observations. Importantly, as all covariates are measured before the cleaning occurs, the matching relies solely on pre-treatment information: the “treatment” itself—whether the cleaner is absent or not—materializes only afterward. Matching on pre-treatment characteristics is widely used in the operations management literature to create comparable treated and control groups in settings such as retail (Jain and Tan 2022) and ride-sharing (Cohen et al. 2021).

We implement one-to-one propensity score matching without replacement with a 0.1 caliper and common support to match each treated cleaning with the closest untreated cleaning according to covariates X_i . These covariates include customer characteristics in cleaning i , such as the number of orders received by the customer up to the focal cleaning

(*CustomerOrderHistory*), the customer's total spending on the platform up to the focal cleaning (*TotalSpending*), and the cost of the focal cleaning (*Price*). Moreover, we also include covariates that are identified as the top three features in our prediction model: these include the total number of cleanings the originally matched cleaner has completed up to the focal cleaning (*CleanerTotalCleaning*), the total number of cleanings they completed in the last 30 days up to the focal cleaning (*CleanerCleaning30Days*), and the total number of no-shows of the cleaner up to the focal cleaning (*CleanerNoShowHistory*). In addition, we exactly match each cleaning with another that took place in the same city for the same product category during the same month-year to further control for product type, local and temporal factors that could potentially affect both no-show events and customer attrition.

We conduct several analyses to verify the quality of our matching algorithm. First, we compute standardized mean differences (SMDs) for all covariates in X_i before and after matching. We observe that all post-matching SMDs are significantly below the commonly accepted threshold of 0.1 (Austin 2009), indicating good balance. Second, t-tests further confirm that the characteristics associated with treated and control cleanings are similar after matching. All t-test p-values exceed 0.05 (minimum p-value = 0.15), indicating no statistically significant differences between treated and untreated cleanings in the matched sample.

When we reestimate our main specification described in Equation 1 on the matched sample, we observe that the effect of no-shows on customer attrition remains positive and significant. Finally, we conduct additional robustness analyses to assess the sensitivity of our results to different matching specifications. These analyses include using an alternative caliper size (0.2 instead of 0.1), and implementing one-to-three matching instead of one-to-one. The results of these alternative specifications remain consistent with our main findings. Collectively, these robustness checks reinforce our main result that customers are significantly less likely to return to the platform following a no-show event.

A.2.2. Alternative Outcome Variables and Specifications. In our main analysis, we examine the impact of no-shows on customer attrition by assessing whether customers place subsequent orders after a cleaning service. Because our dataset ends in February 2020, right-censoring may occur, and customers intending to return later could be misclassified as having churned. To address this issue, we supplement our main analysis with a series of robustness analyses that indicate that our findings remain consistent when using alternative specifications and customer attrition measurements.

First, as an additional robustness check, we implement survival analysis, a methodology widely applied in customer retention studies due to its capacity to handle right-censored data. Specifically, we estimate an accelerated failure time (AFT) model. The AFT model is a flexible model that does not assume proportional hazard rates (Kalbfleisch and Prentice 2002). The AFT model assumes that the survival function $S(t|X)$ at time t is a function of covariates X . The model is specified as follows: $\log(t_i) = (\beta X_i + \epsilon_i)$ where t_i denotes, for the customer who requested the cleaning i , the number of days that have elapsed since the customer's first cleaning on the platform. Similar to our main specification in equation 1, X_i includes the following covariates: *NoShow*, *DiscountAmount*, *CustomerOrderHistory*, *Price*, month-year $\times$ city fixed effects, and product-category fixed effects. We estimate a Weibull AFT model. Our results remain robust under alternative distributions such as the lognormal. In this model, the coefficient β quantifies the effect of variables on the customer survival time on the platform. Consequently, a significant negative coefficient for *NoShow* indicates that experiencing a no-show is associated with shorter observed customer tenure on the platform.

In this analysis, we observe that the coefficient for *NoShow* is negative and significant. This finding is in line with our main result and provides additional support for the conclusion that no-shows increase the likelihood of customer attrition.

Second, we conduct two additional robustness analyses to further assess robustness to right-censoring. In the first approach, we exclude cleanings that took place in the last 90 days of our dataset so that every customer has a minimum of 90 days (three months) to return after a cleaning while maintaining our main dependent variable. This length of inactivity period, and even shorter periods such as 1 month and 1 week, are common benchmarks in both industry reports and academic literature when measuring customer attrition (Neslin et al. 2006; Lemmens and Gupta 2020). Notably, our data indicate that the 90th percentile of the time between consecutive cleanings among customers who return is 78 days. Hence, a 90-day window is a reasonable cutoff in our setting. We estimate equation 1 after excluding the last three months of our dataset. Our analysis shows that the coefficient for *NoShow* is positive and statistically significant, in line with our main findings.

In our third approach, we modify our definition of customer attrition. We define $R90_i$, a binary variable that equals 1 if cleaning i occurs before the final 90 days of our dataset and the customer does not receive any subsequent service within 90 days, and 0 otherwise. Consistent with the previous analysis, we exclude cleanings that occur in the final 90 days of our observation period.² We reestimate equation 1 with this alternative dependent variable while including the same set of control variables described previously. The result of this analysis aligns with our main finding: no-show events significantly increase the odds of customer attrition within a 90-day window.

Overall, all three analyses show that no-show events significantly increase customer attrition. These analyses provide further support that our findings cannot be explained by right-censoring.

A.2.3. Placebo Test. Following Staats et al. (2016), Song et al. (2017), and Jain and Tan (2022), we conduct placebo tests to ensure that the structure of our dataset or model misspecification does not drive our findings. For this analysis, we randomly assign placebo no-show events to different cleanings and reestimate the specification in Equation (1). We repeat this randomization and analysis 1000 times.

In these placebo analyses, we find that 962 out of the 1000 placebo coefficients are not statistically significant at the 5% level. The placebo estimates are symmetric around zero (mean placebo estimate is -0.0008). Importantly, none of the placebo estimates in our analysis is larger than our estimated main effect. This analysis provides reassurance that neither our model nor the structure of our dataset explains the observed effects.

A.2.4. Alternative Mechanisms: Disintermediation. Disintermediation—where service providers and customers bypass the platform and transact offline—remains a challenge for many service businesses. Our platform withholds contact information and disallows messaging between cleaners and customers to impede communication between cleaners and customers on the app. However, disintermediation could still occur if the parties meet in person.

A specific concern for our study is that some instances recorded as “no-shows” might actually represent cases where cleaners arrived at customers’ locations but then proposed off-platform transactions.³ If such behavior were prevalent, it could simultaneously result in an apparent no-show and subsequent customer attrition, potentially confounding our results.

To explore this, we conduct two robustness analyses. First, we control for cleaners’ no-show histories in the main specification described in Equation 1. The rationale is that if certain cleaners systematically engage in disintermediation that is masked as no-shows, their historical record would control for this tendency. The result of this analysis

is very similar to our main analysis.⁴ This finding indicates that individual cleaner propensities to use no-shows as a disintermediation strategy do not drive our results.

Second, we restrict our analysis to no-show incidents that were followed by same-day restoration—that is, we only retain no-shows for which the customer received the cleaning service through the platform on the same day as the initial no-show and exclude all other no-show instances.⁵ This approach aims to isolate no-show events where disintermediation is unlikely to have occurred, since customers who receive off-platform service are unlikely to reschedule and receive the cleaning through the platform on the same day. For this subsample, we estimate the model described in Equation 1. We find that the estimate for no-show is very similar to the main specification, further supporting the conclusion that the observed relationship between no-shows and customer attrition cannot be explained by disintermediation.

A.2.5. Details on Prediction Model. In this section, we discuss the features that are included in our prediction model and the details of model selection.

For the cleaner, we measure activity on the platform by tracking total cleanings completed before the current cleaning. In addition, we separately compute this metric for the past 7, 15, and 30 days. We also measure the cleaner’s no-show history (i.e., the total number of no-shows before the focal cleaning) and the fatigue level (measured as the number of days since the last service the cleaner completed), and the geographical distance between the customer and the cleaner.⁶

For the customer, we include the total number of orders before the current cleaning, tenure on the platform (measured by the number of days since their first cleaning on the platform), and total spending before the current cleaning.

For the cleaning itself, we include the scheduled start hour, urgency (measured by the days between the customer’s order date and the requested service date), month-year $\times$ city controls for the month-year $\times$ city of the service, product category, whether any extra services are requested, extra service types, whether there are any pets in the house, price, discount amount, and weather conditions on the day of service including temperature, precipitation, snow depth, wind speed, and visibility.

Model Selection. In our sample, 3.51% of the observations are no-shows. We estimate gradient-boosted trees because they constitute a state-of-the-art prediction method due to their outstanding performance on sparse datasets.

As our outcome variable is binary, we define the objective function as a logistic loss function. We set the learning rate η equal to 0.1 and use an 80/20 training/testing split. To select the optimal hyperparameters—including maximum tree depth, minimum child weight, column ratios, row subsample ratios, partition penalty, and weight penalty—we perform a grid search over the set $\{3, 4, 5, 6\} \times \{1, 2\} \times \{0.6, 0.8, 1\} \times \{0.8, 1\} \times \{0, 4, 8\} \times \{0, 1, 2\}$ using 10-fold cross-validation.

For each fold, we identify the hyperparameter vector that minimizes the out-of-sample loss function value. We then select the final hyperparameter vector corresponding to the global minimum loss function value across all folds. After determining the best hyperparameter values, we retrain the model on the full training set (80% of the data) and evaluate its predictive performance on the holdout set (20%).

Given the sparsity of no-show events, we evaluate model performance using the F1 score. The F1 score, which captures the performance of both precision and recall, gives a better and more transparent picture of the prediction model’s performance in sparse datasets. In contrast, other metrics not accounting for data imbalance might give misleading performance scores (Jeni et al. 2013).

Results. The F1 score of our prediction model is 0.27, indicating limited predictive performance. The most predictive features are cleaner related. Specifically, the top three features are the cleaner’s no-show history (positively correlated with no-shows), the total number of cleanings the focal cleaner completed before the focal cleaning (negatively correlated with no-shows), and the number of cleanings the cleaner completed in the past 30 days (negatively correlated with no-shows). Although these three features provide the most signal, the overall performance of the model is poor: an F1 score of 0.27 is much lower than the rule-of-thumb threshold of 0.7, the threshold which indicates strong predictive performance. This result underscores the inherent difficulty of predicting no-shows *ex ante*.

Appendix B: Details of the Customer Lifetime Value Model

In this section, we first provide an extensive discussion about the extended pnb model for purchase and attrition, the Gamma-Gamma model for spending. Next, we provide details about the model implementation, model selection, and goodness-of-fit checks.

Repeat Purchase and Attrition Model. To model customers’ repeat purchases and attrition, we use an extension of the Pareto/negative binomial distribution model (PNBD). PNBD is the most widely used framework among researchers and practitioners in *non-subscription* settings like ours (Schmittlein et al. 1987, Oblander and McCarthy 2021). In PNBD, customers can be in one of the two states: “Alive” or “Dead” (churn). When alive, customer i ’s purchases follow a Poisson process with rate λ_i ; she stays alive for a lifetime whose length follows an exponential distribution with rate μ_i . λ_i and μ_i follow two different and independent Gamma distributions.

We implement an extended PNBD framework developed by Bachmann et al. (2021) that incorporates time-varying variables in the *rates* of the Poisson purchase process and the exponentially distributed lifetime. Specifically, customer i ’s purchases (cleaning requests in our case) follow a nonhomogeneous Poisson process with time-varying rate $\lambda_i(t) = \lambda_0 \exp(\gamma'_p \mathbf{X}_{it}^P)$, and her lifetime follows an exponential distribution with rate $\mu_i(t) = \mu_0 \exp(\gamma'_a \mathbf{X}_{it}^A)$. The covariate vectors $\mathbf{X}_{it}^P$ can include time-varying variables. The baseline purchase and attrition rates are independently distributed according to gamma distributions: $\lambda_0 \sim \Gamma(r, \alpha)$, $\mu_0 \sim \Gamma(s, \beta)$. The individual likelihood is then constructed as the product of two components: the density of the inter-purchase times and a survivor function that evaluates the customer’s probability of being alive beyond their last observed transaction. The parameters $\Theta := (r, \alpha, s, \beta, \gamma_a, \gamma_p)$ are estimated by maximizing the likelihood across all customers. Estimation of the purchase process parameters relies on variation in inter-purchase times within and across customers, while attrition parameters are estimated from variation in recency (i.e., time from last purchase to end of observation), conditional on purchase history.

This extended model suits our purposes because our covariate of interest—customer i ’s no-show incidents—can vary over time. We include whether customer i experienced a no-show by time t in both $\mathbf{X}_{it}^P$ and $\mathbf{X}_{it}^A$. Following Bachmann et al. (2021) and references therein, our main specification further includes seasonality at time t in $\mathbf{X}_{it}^P$ to capture aggregate temporal shocks that may affect purchasing behavior. We also incorporate customer i ’s cumulative spending by time t into the attrition process.⁷ With the estimated parameters, two key metrics are derived from the model: 1) the probability of a customer being alive after time T given her past purchase history, and 2) if alive after time T , the expected number of purchases made between T and a future time $T' > T$. These two components together determine the customer’s conditional expected number of purchases during (T, T') .

Spending Model. We estimate the mean spending per transaction, M , for each customer through the Gamma-Gamma model proposed by Fader et al. (2005) and Colombo and Jiang (1999). The spending model is estimated

independently from the extended PNBD model. For each customer i , her spending per transaction, $\mathbf{z}_i$ are assumed to be i.i.d. gamma random variables with shape parameter p and scale parameter ν . To capture customer heterogeneity in the average transaction value across customers, the model further assumes ν follows a Gamma distribution, $\nu \sim \Gamma(q, \pi)$. Given that customer i 's observed average spending per transaction is $m_x = \sum_{k=1}^x z_{i,k} / x$ over x transactions, we can derive the density function of m_x for each customer. The parameters (p, q, π) are estimated by maximizing the likelihood function, defined as the product of individual densities of m_x . With the estimated parameters, we can then derive the posterior distribution of ν , and compute the expected mean spending per transaction conditional on individuals' observed x and m_x , $\mathbb{E}[M|m_x, x]$. Note that we assume customers' average spending per transaction is not affected by no-show incidents.

The model-predicted CLV during (T, T') for each customer is then calculated as the product of two terms: the conditional expected number of purchases during (T, T') estimated from the purchase and attrition model and the conditional expected mean spending per transaction from the spending model.

Estimation and Goodness-of-fit. Following the literature (Bachmann et al. 2021), we aggregate our data to customer-week level, where t denotes week t . We implement both models using the *pnb*d function in R package "CLVTools" (Bachmann et al. 2022). We calibrate our models with data from the first 52 weeks (i.e., one year, $T = 52$) (as indicated by the dashed line in Figure B.1), and predict over the remaining 8 holdout weeks (i.e., end of our observed data, $T' = 60$). Estimation results of this model are presented in Table B.1.

Parameter	r	α	s	β	$\gamma_{a, \text{HadNoShow}}$	$\gamma_{a, \text{SpendingSoFar}}$	$\gamma_{p, \text{HadNoShow}}$	$\gamma_{p, \text{Seasonality}}$
Estimate	0.363***	5.971***	0.880***	28.834***	0.229 [†]	-0.138***	-0.109**	0.095***
(s.e.)	(0.009)	(0.214)	(0.143)	(6.638)	(0.127)	(0.028)	(0.035)	(0.007)

Table B.1 Extended PNBD Model Estimates

This table reports the estimates solved from the extended pnb model using the *pnb*d function from the *CLVTools* R package (Bachmann et al. 2022). Parameters r , α , s , and β control the gamma distributions of interpurchase times and dropout risks. The γ parameters represent time-varying covariates influencing purchasing behavior and attrition risk. Model selected based on BIC as in Bachmann et al. (2021). *** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$, [†] $p < 0.1$.

Before proceeding to the counterfactual analysis, we first assess the model's predictive accuracy. Following Bachmann et al. (2021), we compare model predictions to actual outcomes for two key metrics: 1) the number of purchases during the holdout period, and 2) total customer spending during the holdout period. Actual spending serves as the realized CLV, against which we compare model-predicted CLV. As our primary goal is to understand how no-shows influence purchase and attrition behavior, and therefore CLV, these two metrics directly capture the outcomes most relevant to our analysis. We visualize the model fit in Figure B.1.

The top-left panel plots the number of observed transactions during the estimation period (first 52 weeks) on the horizontal axis. Customers are grouped by their total number of transactions (that is, 1, 2, ..., 10+). The vertical axis shows the average number of transactions during the holdout period for each group, with separate points and curves for actual (black) and model-predicted (grey) values. The top-right panel follows the same logic but compares spending. The horizontal axis shows each customer's realized CLV during the estimation period, and the vertical axis shows CLV during the holdout period. The close alignment between actual and predicted values suggests that the model captures key patterns in customer spending. Furthermore, the correlation between actual and predicted holdout-period

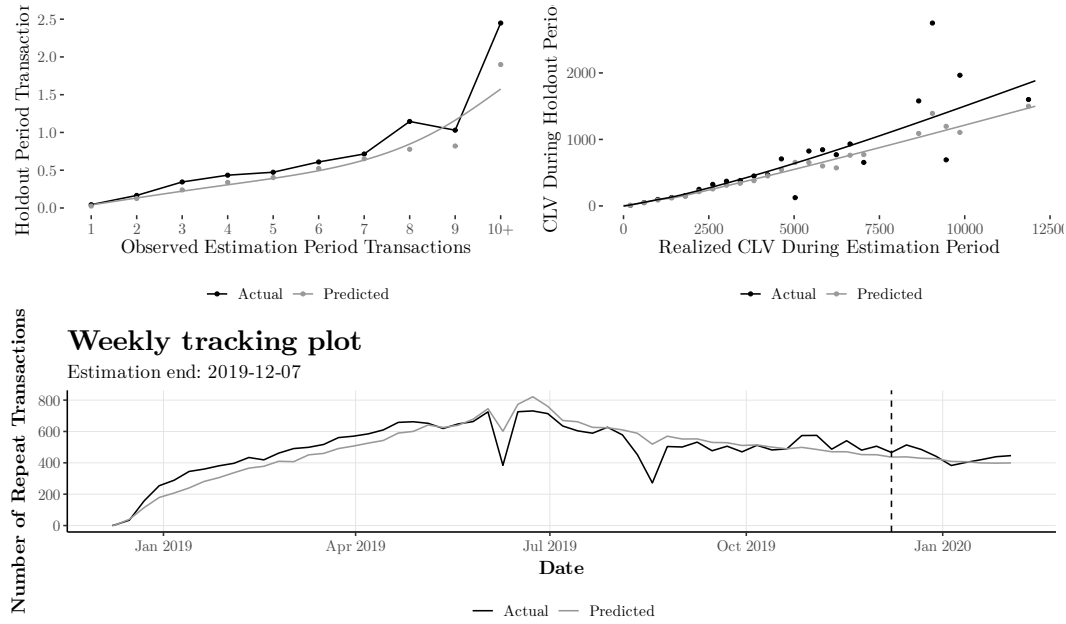


Figure B.1 CLV Model Fitting Diagnosis

This figure plots the observed (black curve) and predicted (grey curve) number of transactions and customer lifetime value to visualize the goodness of fit of our CLV model. Note that to make sure each customer has enough transaction history for modeling, we include only customers whose first cleaning happened before June 27, 2019 (half a year after beginning of the dataset in Dec 2018) in the CLV modeling, which contribute a total of 43,278 transactions. To construct this figure, we ignore the customers' order history prior to December 2018. We follow the examples in [Bachmann et al. \(2021\)](#) and aggregate our data to customer-week level for this analysis. The dashed line indicates the start of the prediction period. This figure suggests that the CLV model fits our data well.

transactions is 0.77, and the correlation between total customer spending and predicted CLV is 0.74. For reference, [Bachmann et al. \(2021\)](#) report correlations around 0.60 in their examples.

To further assess model performance at the platform level, we aggregate individual predictions across all customers. Platform-level model fit is supported by the bottom panel of Figure B.1, which tracks weekly repeat purchases—transactions made by returning customers. The predicted curve (grey) closely follows the actual trend (black), including seasonal fluctuations, indicating the model effectively captures both customer churn and purchase frequency over time. We refer back to Section 4.4 for the discussion of the model estimates and the counterfactual analysis.⁸

B.1. CLV Model Covariates

We estimate eight CLV models with different sets of covariates. The main specification is selected among these eight candidate sets based on BIC following [Bachmann et al. \(2021\)](#). These covariate sets are $\mathbf{X}_{it}^A \times \mathbf{X}_{it}^P \in \left\{ \{HadNoShow_{i,t}\}, \{HadNoShow_{i,t}, SpendingSoFar_{i,t}\}, \{HadNoShow_{i,t}, Seasonality_t\}, \{HadNoShow_{i,t}, SpendingSoFar_{i,t}, Seasonality_t\} \right\} \times \left\{ \{HadNoShow_{i,t}\}, \{HadNoShow_{i,t}, Seasonality_t\} \right\}$.

$HadNoShow_{i,t}$ equals 1 if customer i had experienced a no-show prior to time t (week t). $SpendingSoFar_{i,t}$ is customer i 's cumulative spending on the platform up to, but not including, week t . $Seasonality_t$ is the platform's total revenue aggregated over all customers in week t , standardized by its standard deviation across the estimation periods.

Our main model, with $\mathbf{X}_{it}^A = \{HadNoShow_{i,t}, SpendingSoFar_{i,t}\}$ and $\mathbf{X}_{it}^P = \{HadNoShow_{i,t}, Seasonality_t\}$, achieves the lowest BIC among all specifications at optimality.

Appendix C: Additional Details for Instrumental Variable Analyses

C.1. First Stages for the Two-Stage Residual Inclusion Approach

	<i>FrustrationDiscountAmount</i> (1) OLS	<i>NoShowFrustrationDiscount</i> (2) OLS	<i>SameDayRestored</i> (3) Logit	<i>OneHourRestored</i> (4) Logit
<i>NoShow</i>	4.820** (1.625)	6.113*** (1.619)	23.798*** (0.074)	16.630*** (0.092)
<i>FDOtherCityDate</i>	0.251*** (0.060)	-0.097*** (0.013)		
<i>NoShow × FDOtherCityDate</i>	1.717*** (0.316)	2.196*** (0.314)		
<i>OccupancyOtherDistricts</i>			-0.013*** (0.002)	-0.025*** (0.006)
Observations	81,959	81,959	81,959	81,959
Log Likelihood	-364,877.8	-287,545.7	-1,771.4	-581.9

Table C.1 First Stages of the IV Approach

This table reports the results of first-stage analysis of 2SRI. All models include *CustomerOrderHistory*, *Price*, *DiscountAmount* controls, as well as month-year $\times$ city fixed effects and product-category fixed effects. Robust standard errors are in parentheses. *** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$, † $p < 0.1$.

C.2. Sensitivity of Frustration Discount IV Estimates to Potential Violations of the Exclusion Restriction

In Section 5.1, we leverage *FDOtherCityDate* as an instrument for *FrustrationDiscountAmount* and find that frustration discounts reduce customer attrition following no-shows. We provide arguments supporting the validity of *FDOtherCityDate* as an instrument, but the exclusion restriction is fundamentally untestable. In this section, to assess the robustness of our IV estimates to potential violations of the exclusion restriction, we run two additional analyses.

Additional Controls. To address the potential concern that date $\times$ city level shocks may simultaneously affect discounting behavior and customer attrition, we introduce additional controls capturing platform dynamics in the same city on the same day of the focal cleaning, following prior literature (Xu et al. 2023). Specifically, in both stages of the IV estimation, we control for: (i) the total number of cleanings scheduled in the focal city on that date, excluding the focal cleaning, and (ii) the total number of those cleanings that received frustration discounts.

By capturing daily demand conditions and the extent of service degradation or promotion issuance among other customers in the same date-city, these controls ensure that the remaining variation in our instruments affects customer attrition only through the focal cleaning's frustration discount—rather than through the broader platform dynamics.

Our IV results remain robust when these controls are included. The relevance still holds, as can be observed from first-stage analysis ($\beta=0.21$, $p<0.001$; $\beta=2.2$, $p<0.001$). The estimated effects from the second stage are similar to the main specification, suggesting that our findings are not driven by these kinds of platform dynamics. The coefficient of *FrustrationDiscountAmount* is ($\beta = -0.044$, $p=0.039$). The total effect of *FrustrationDiscountAmount* + *NoShow* $\times$ *FrustrationDiscountAmount* is also negative and statistically significant ($\beta = -0.019$, $p=0.072$). In addition, we compute bootstrap standard errors clustered at the date $\times$ city level, and the results remain robust under bootstrap standard errors as well.

Conley Bounds. To assess the sensitivity of our IV estimates to potential violations of the exclusion restriction, we follow Conley et al. (2012) and conduct a plausibly-exogenous (Conley-style) sensitivity analysis using the union-of-confidence-intervals (UCI) approach. Conley bounds are commonly used to evaluate the robustness of IV estimates to

violations of the exclusion restriction assumption in prior studies (e.g., [Frake and Harmon 2023](#); [Marinoni and Roche 2025](#); [Guo et al. 2025](#)).

In this approach, we allow the instruments to have direct effects on the outcome (customer attrition), rather than assuming that the exclusion restriction holds perfectly. We first estimate the reduced-form relationship between the instruments and attrition while controlling for the full set of covariates and fixed effects used in the IV specification. We then allow the instruments' direct effects to vary as a fraction of the reduced-form magnitudes. For each assumed violation level, we re-estimate the second-stage model and compute the corresponding confidence interval, and report the union of these intervals across all values considered, following the union-of-confidence-intervals procedure of [Conley et al. \(2012\)](#), which provides conservative inference without imposing distributional assumptions on the violation size.

Table C.2 reports the resulting UCI 90% confidence intervals under increasing degrees of exclusion restriction violations. Across all values considered, the bounds for *FrustrationDiscountAmount* remain negative and exclude zero, confirming that larger frustration discounts reduce customer attrition. For the combined effect (*FrustrationDiscountAmount* + *NoShow* × *FrustrationDiscountAmount*), the bounds remain negative and exclude zero for up to 15% of the reduced-form magnitude, indicating that our conclusion—that issuing larger frustration discounts mitigates customer attrition following no-shows—is robust to plausible deviations from strict instrument exogeneity.

	Violation size (as % of reduced form)	Lower Bound	Upper Bound
<i>FrustrationDiscountAmount</i>	5%	−0.0786	−0.0178
	10%	−0.0810	−0.0154
	15%	−0.0833	−0.0131
	20%	−0.0857	−0.0107
<i>FrustrationDiscountAmount</i> + <i>NoShow</i> × <i>FrustrationDiscountAmount</i>	5%	−0.0371	−0.0021
	10%	−0.0381	−0.0011
	15%	−0.0391	−0.00015
	20%	−0.0400	0.00081

Table C.2 Conley Bounds for Frustration Discount IV Estimates

This table reports the union of 90% confidence intervals for coefficients from the second-stage logistic regression under varying degrees of exclusion restriction violation.

Appendix D: Proof of Proposition 1

We rewrite the objective function as follows.

$$\mathbb{E}_{C|\overline{C}} \left[\mathbb{E}_{\alpha} [R(\alpha, C, d_H, d_L | \theta, \psi_{\tau_c}, \mu) | C] \right] - \mu \overline{C}, \text{ where } \theta = (L, v_H, v_L, p, m, m_o, \beta_H, \beta_L). \quad (8)$$

Proof of Part 1. We start by simplifying the maximization problem. Note first, since $(1 - \gamma)\alpha L - z \geq 0$, and $\gamma\alpha L - y \geq 0$, the optimal discount $d_{\tau_c}^*$ that solves for $\max_{d_{\tau_c} \geq 0} s_{\tau_c}(d_{\tau_c})$ will also solve for the original objective function $d_{\tau_c}^*$. Therefore, $d_{\tau_c}^* = \arg \max_{d_{\tau_c} \geq 0} s_{\tau_c}(d_{\tau_c})$. Next, given $k_{o,\tau_c} \geq s_{\tau_c}(d_{\tau_c}^*)$, we conclude $z^* = \min\{(1 - \gamma)\alpha L, C - y\}$. We further denote $g_{\tau_c}^* := k_{o,\tau_c} - s_{\tau_c}(d_{\tau_c}^*)$. Next given α and C , consider

$$\max_{0 \leq y \leq \min\{\gamma\alpha L, C\}} H := \{(k_{o,H} - s_H(d_H^*))y + (k_{o,L} - s_L(d_L^*)) \min\{(1 - \gamma)\alpha L, C - y\}\}.$$

We conclude $y^* = \min\{\gamma\alpha L, C\}$ when $g_H^* \geq g_L^*$. The solution when $g_H^* < g_L^*$ is included in Appendix F. Substituting $s_{\tau_c}(d_{\tau_c}^*)$, y^* , and z^* into the optimization problem, we have,

$$\max_{\bar{C} \in \mathbb{Z}^+ \cup \{0\}} \mathbb{E}_{C|\bar{C}} \left[\mathbb{E}_\alpha \left[k_H(1-\alpha)\gamma L + k_L(1-\alpha)(1-\gamma)L + \right. \right. \\ \left. \left. k_{o,H} \min\{\gamma\alpha L, C\} + k_{o,L} \min\{(1-\gamma)\alpha L, (C - \gamma\alpha L)^+\} + \right. \right. \\ \left. \left. (\gamma\alpha L - \min\{\gamma\alpha L, C\})s_H(d_H^*) + ((1-\gamma)\alpha L - \min\{(1-\gamma)\alpha L, (C - \gamma\alpha L)^+\})s_L(d_L^*) \middle| C \right] \right] - \mu\bar{C}.$$

The notation $C|\bar{C}$ indicates that C is a random variable that follows $\text{Bin}(\bar{C}, \beta_s)$. Note that $\min\{(1-\gamma)\alpha L, (C - \gamma\alpha L)^+\} = \min\{\alpha L, C\} - \min\{\gamma\alpha L, C\}$. Conditioning on C , we start by simplifying the inner expectation on α .

$$\mathbb{E}_\alpha \left[\min\{\gamma\alpha L, C\} (k_{o,H} - k_{o,L} - s_H(d_H^*) + s_L(d_L^*)) + \min\{\alpha L, C\} (k_{o,L} - s_L(d_L^*)) \right] + \\ (\gamma\mu_\alpha L(s_H(d_H^*) - k_H) + (1-\gamma)\mu_\alpha L(s_L(d_L^*) - k_L) + (k_H\gamma L + k_L(1-\gamma)L)),$$

where $\mu_\alpha = \mathbb{E}_\alpha[\alpha]$. We write the above as $\mathbb{E}_\alpha \left[\min\{\gamma\alpha L, C\} M_1 + \min\{\alpha L, C\} M_2 \right] + M_3$. Note that M_i are parameters, i.e., not functions of any random variables. We now estimate the expectation over C , where $C \sim \text{Bin}(\bar{C}, \beta_s)$.

$$\mathbb{E}_{C|\bar{C}} [\mathbb{E}_\alpha \min\{\gamma\alpha L, C\} | C] = \beta_s \bar{C} - \gamma L \mathbb{E}_{C|\bar{C}} \left[\int_0^{C/\gamma L} F_\alpha(x) dx \right] = \beta_s \bar{C} - \gamma L \mathbb{E}_{C|\bar{C}} \left[\int_0^\infty F_\alpha(x) \mathbb{I}\{x \leq C/\gamma L\} dx \right] \\ = \beta_s \bar{C} - \gamma L \int_0^\infty F_\alpha(x) \mathbb{P}(C \geq \gamma L x) dx$$

The first equality follows from $\mathbb{E}_\alpha[\min\{\gamma\alpha L, C\}] = C - \gamma L \int_0^{C/\gamma L} F_\alpha(x) dx$. The third equality follows from Fubini's theorem. $|F_\alpha(x)| < 1$ since F_α is the c.d.f. of α . Further, $\mathbb{P}(C \geq \gamma L x) = 0$ for any $x > \bar{C}/(\gamma L)$. Therefore, $\int_0^\infty |F_\alpha(x)| \mathbb{P}(C \geq \gamma L x) dx$ is finite. We can therefore write the expectation term as an iterated integral and change the order of integration, i.e., interchange the expectation and the integral. $\mathbb{E}_{C|\bar{C}} [\mathbb{E}_\alpha \min\{\alpha L, C\} | C]$ can be similarly estimated as above.

We use trapezoidal approximation to numerically calculate the integral in $\mathbb{E}_{C|\bar{C}} [\mathbb{E}_\alpha \min\{\gamma\alpha L, C\} | C]$. Let $h_{C|\bar{C}}(y) = \mathbb{P}(C \geq y) = \sum_{i=y}^{\bar{C}} \mathbb{P}(C = i)$, evaluated at integer grid points. We have,

$$\mathbb{E}_{C|\bar{C}} [\mathbb{E}_\alpha \min\{\gamma\alpha L, C\} | C] = \beta_s \bar{C} - \gamma L \int_0^\infty F_\alpha(x) \mathbb{P}(C \geq \gamma L x) dx = \beta_s \bar{C} - \int_0^\infty F_\alpha((\gamma L)^{-1} y) \mathbb{P}(C \geq y) dy \\ \approx \beta_s \bar{C} - \frac{1}{2} \left[\sum_{y=0}^{\bar{C}} 2F_\alpha((\gamma L)^{-1}(y)) h_{C|\bar{C}}(y) - F_\alpha(0) h_{C|\bar{C}}(0) - F_\alpha((\gamma L)^{-1}(\bar{C})) h_{C|\bar{C}}(\bar{C}) \right].$$

The second equality follows from change of variables, and the last approximate equality follows from trapezoidal approximation. We can numerically integrate $\mathbb{E}_{C|\bar{C}} [\mathbb{E}_\alpha \min\{\alpha L, C\} | C]$ by similar arguments. We can therefore write down the optimality condition for $\bar{C}^*$,

$$\mathbb{E}_{C|\bar{C}} \left[\mathbb{E}_\alpha [\min(\gamma\alpha L, C) M_1 + \min(\alpha L, C) M_2] | C \right] - \mu\bar{C} \geq \mathbb{E}_{C|\bar{C}'} \left[\mathbb{E}_\alpha [\min(\gamma\alpha L, C) M_1 + \min(\alpha L, C) M_2] | C \right] - \mu\bar{C}',$$

for all $\bar{C}' \neq \bar{C}$. This proves the expression of $\bar{C}^*$ and completes the proof of part 1.

LEMMA 1. *The objective function $\mathbb{E}_{C|\bar{C}} [\mathbb{E}_\alpha [\min\{\gamma\alpha L, C\} M_1 + \min\{\alpha L, C\} M_2] | C] + M_3 - \mu\bar{C}$ is (discrete) concave in $\bar{C}$ and unimodal.*

See Page 11 for proof of Lemma 1. By Lemma 1, the objective function is unimodal with the integer constraints on $\bar{C}$. Therefore, the sufficient condition for $\bar{C}^*$ to be optimal is when the above inequality is true at $\bar{C}' = \bar{C}^* + 1$ and $\bar{C}' = \bar{C}^* - 1$. We apply this sufficient condition for an efficient numerical analysis.

Proof of Part 2. We omit the subscript τ_c for ease of notation. The value recovered through a discount d is, $s(d; \psi) = v(1 - \beta(d; \psi)) - g(d)$. Its derivative is $s'(d) = -v\beta'(d; \psi) - g'(d)$. Recall s is concave in d . The optimal discount $d^* = 0$ if $s'(d) \leq 0$ for all $d \in [0, \infty)$. By concavity, $s''(d) \leq 0$, so $s'(0) \geq s'(d)$ for all $d \geq 0$. It therefore suffices to check if $s'(0) \leq 0$. That is, the marginal benefit of issuing the first unit of discount does not cover its marginal cost. Defining $\Psi = \{\psi \mid -v\beta'(0; \psi) \leq g'(0)\}$, we have $d^* = 0$ for all $\psi \in \Psi$.

If the platform chooses to have $\bar{C} = 0$, the optimal profit from a discount-only policy is,

$$k_H(1 - \mu_\alpha)\gamma L + k_L(1 - \mu_\alpha)(1 - \gamma)L + \gamma\mu_\alpha L s_H(d_H^*) + (1 - \gamma)\mu_\alpha L s_L(d_L^*).$$

The difference in optimal profits between a standby cleaner plus discounts hybrid strategy and a discount-only strategy is then,

$$\max_{\bar{C} \in \mathbb{Z}^+ \cup \{0\}} \mathbb{E}_{C|\bar{C}} \left[\mathbb{E}_\alpha \left[(k_{o,H} - s_H(d_H^*)) \min\{\gamma\alpha L, C\} + (k_{o,L} - s_L(d_L^*)) \min\{(1 - \gamma)\alpha L, (C - \gamma\alpha L)^+\} \right] | C \right] - \mu\bar{C}.$$

By condition, $s_{\tau_c}(d_{\tau_c})$ (weakly) increases in ψ_{τ_c} for all $d \geq 0$, therefore $s_{\tau_c}(d_{\tau_c}^*)$ (weakly) increases in ψ_{τ_c} . Hence, given any $\bar{C}$ and ψ_L , the expectation term in the objective function decreases in ψ_H . In the proof of Lemma 1, we have shown that given ψ , the expectation term is nondecreasing and concave. Denote the expectation term as $R(\bar{C}, \psi_H)$, we start by showing that $\bar{C}^*(\psi_H) = \arg \max_{\bar{C}} R(\bar{C}, \psi_H) - \mu\bar{C}$ decreases in ψ_H . Let $\psi_H = -\tilde{\psi}_H$. Note the coefficient $k_{o,H} - s_H(d_H^*)$ decreases as ψ_H increases. Therefore, $R(\bar{C}, \tilde{\psi}_H) - \mu\bar{C}$ has increasing differences in $(\bar{C}, \tilde{\psi}_H)$. Intuitively, when the high-value customers become more sensitive to discounts, the benefits of adding additional standby cleaners decrease. Following Milgrom and Shannon (1994), we have $\bar{C}^*$ increases in $\tilde{\psi}_H$, therefore decreases in ψ_H . Since $R(\bar{C}, \psi_H)$ is bounded, when μ is sufficiently large, following from $\bar{C}^*(\psi_H)$ decreases in ψ_H , there exists a $\psi^0(\mu)$ such that, when $\psi_H > \psi^0(\mu)$, $\bar{C}^* = 0$. The same logic applies when ψ_H is fixed, and ψ_L varies. This completes the proof of part 2.

Proof of Lemma 1. First, M_3 is irrelevant to the optimality of $\bar{C}$. We only need to consider the concavity of $\mathbb{E}_{C|\bar{C}} \left[\mathbb{E}_\alpha \left[\min(\gamma\alpha L, C)M_1 + \min(\alpha L, C)M_2 \right] | C \right] - \mu\bar{C}$. Let $G(\bar{C}) = \mathbb{E}_{C|\bar{C}} \left[\mathbb{E}_\alpha \left[\min(\gamma\alpha L, C)M_1 + \min(\alpha L, C)M_2 \right] | C \right] = \mathbb{E}_{C|\bar{C}} \left[\mathbb{E}_\alpha g(\alpha, C) | C \right]$. Note first, given the conditions, we have $M_1(\cdot), M_2(\cdot) > 0$, and further we have seen that $M_1(\cdot), M_2(\cdot)$ are independent of α and C . It is easy to see that given α , $g(\alpha, C)$ is the sum of two nondecreasing and concave functions in C for all $\alpha \in (0, 1)$. Therefore, $g(\alpha, C)$ is also nondecreasing and concave in C . $\mathbb{E}_\alpha g(\alpha, C) = \int_\alpha g(\alpha, C) f(\alpha) d\alpha$ is then also concave in C since $f(\alpha) > 0$ and $g(\alpha, C)$ is concave in C at each α (Boyd and Vandenberghe 2004, p.79). Lastly, taking the expectation over $C \sim \text{Bin}(\bar{C}, \beta_s)$. We denote $f(C) := \mathbb{E}_\alpha g(\alpha, C)$. $\mathbb{E}_{C|\bar{C}} f(C) = \sum_{k=0}^{\bar{C}} f(k) \mathbb{P}(k)$. We want to show $\mathbb{E}_{C|\bar{C}} f(C)$ is concave in $\bar{C}$ when $f(C)$ is concave and C follows a binomial distribution.

Let $C_+ \sim \text{Bin}(\bar{C} + 1, \beta_s)$, $C \sim \text{Bin}(\bar{C}, \beta_s)$, and $C_- \sim \text{Bin}(\bar{C} - 1, \beta_s)$. First note that $C_+ = C + B$, where $B \sim \text{Ber}(\beta_s)$. Therefore, $\mathbb{E}_{C|\bar{C}_+} f(C_+) = \mathbb{E}_{C|\bar{C}} [\mathbb{E}_B [f(C + B) | B]] = \beta_s \mathbb{E}_{C|\bar{C}} f(C + 1) + (1 - \beta_s) \mathbb{E}_{C|\bar{C}} f(C)$. This equivalence can also be seen from expanding the binomial coefficients and applying Pascal's rule.

To prove concavity of $\mathbb{E}_{C|\bar{C}}f(C)$ in $\bar{C}$, it suffices to show $\Delta_f(\bar{C}) = \mathbb{E}_{C|\bar{C}_+}f(C_+) - \mathbb{E}_{C|\bar{C}}f(C)$ decreases in $\bar{C}$, i.e., the second difference is nonpositive (Ninh et al. 2020). Simplifying Δ_f , we have $\Delta_f(\bar{C}) = \beta_s(\mathbb{E}_{C|\bar{C}}f(C+1) - \mathbb{E}_{C|\bar{C}}f(C)) = \beta_s(\mathbb{E}_{C|\bar{C}}[f(C+1) - f(C)])$.

First, we have proved above that $f(k)$ is concave. Therefore, $h(k) := f(k+1) - f(k)$ decreases in k . Second, suppose $C_1 \sim \text{Bin}(\bar{C}_1, \beta_s)$ and $C_2 \sim \text{Bin}(\bar{C}_2, \beta_s)$, where $\bar{C}_1 > \bar{C}_2$. Since C_1 can be again decomposed as C_2 plus some additional independent Bernoulli trials, it is easy to see that $\mathbb{P}(C_1 \geq k) \geq \mathbb{P}(C_2 \geq k)$, for all $k \in \mathbb{R}^+$. Therefore, C_1 first order stochastically dominates C_2 , that is, $C_1 \geq_{st} C_2$, by definition. Given $h(\cdot)$ is decreasing, and the stochastic dominance relationship, we have $\mathbb{E}_{C_1|\bar{C}_1}[h(C_1)] \leq \mathbb{E}_{C_2|\bar{C}_2}[h(C_2)]$ (Shaked and Shanthikumar 2007, p.4). Since $\Delta_f = \beta_s \mathbb{E}_{C|\bar{C}}[h(C)]$, and $\beta_s > 0$, we conclude Δ_f decreases in $\bar{C}$. We therefore show that $\mathbb{E}_{C|\bar{C}}f(C)$ is concave in $\bar{C}$. We have therefore shown that $G(\bar{C})$ is nondecreasing and concave in $\bar{C}$. Given $-\mu\bar{C}$ is linear and decreasing, we conclude that $G(\bar{C}) - \mu\bar{C}$ is (discrete) concave in $\bar{C}$ and unimodal. This completes the proof.

In the numerical studies, we assume $s(d) = v(1 - \beta \exp(-\psi d)) - d$. Solving the inequality $s'(0) \leq 0$, we have $\psi \leq \frac{1}{v\beta}$. That is, the platform is better off not offering a discount when $\psi \leq \frac{1}{v\beta}$. We also assume $\alpha \sim \text{Unif}$ in the numerical studies.

Appendix E: Model Parameters Calibrated From Data

We focus on one major city from our dataset. The median number of daily cleaning requests in this city is $L = 108$. On average, daily mean no-show rate is $\mu_\alpha = 0.035$, with a standard deviation of $\sigma_\alpha = 0.019$. We assume α follows a uniform distribution, i.e. $\alpha \sim \text{Unif}$. On average, a customer pays the platform $p = 121$ LC per cleaning, while the platform pays the cleaner $m = 111$ LC per cleaning.⁹ We assume that if a standby cleaner is dispatched, the platform pays a 5% premium, $m_o = 5$ LC (roughly 0.05×111). We estimate the distribution of CLVs based on our CLV analysis.¹⁰ Customers are then categorized into high-value and low-value customers based on whether their CLV is above or below the distribution mean. The median CLV among high-value customers is $v_H = 460$, roughly 4 cleanings, while the median CLV among low-value customers is $v_L = 90$, roughly 1 cleaning. We assume $\beta_{\tau_c}(d; \psi_{\tau_c}) = \beta_{\tau_c} \exp(-\psi_{\tau_c} d)$, and $g(d) = d$. Using empirical estimates (see Column (2) in Table 4) and the base attrition probabilities in low-value and high-value customers, it can be calculated that no-shows increase attrition probability by 44% for high-value customers and by 7% for low-value customers. Using the coefficients of FrustrationDiscountAmount and NoShowFrustrationDiscount in Column 3 of Table 5, we can also estimate the customers' responsiveness to the frustration discount in the face of no-show to be $\psi_L = \psi_H = 0.02$.

Optimal Strategy for the Focal Platform. We assume the platform pays $\mu = 15$ LC per standby cleaner, who is available with probability $\beta_s = 0.8$ at this price. There are roughly $\gamma = 20\%$ high-value customers. The optimal strategy for our focal platform is to retain 7 standby cleaners in this city. When no-shows occur, high-value customers are prioritized for recovery. All available standby cleaners should be fully utilized. For high-value customers whose cleanings cannot be recovered, the platform should offer up to 65 LC in frustration discounts (equivalent to a 54% discount on one cleaning). No discounts should be extended to low-value customers. Maintaining 7 standby cleaners means reserving a standby pool corresponding to approximately 6% ($=7/108$) of daily cleaning volume to enable prompt restoration in the event of a no-show. Importantly, in our model, the optimal number of standby cleaners is a function of platform scale. Accordingly, the magnitude of 7 standby cleaners should be interpreted relative to the size of daily cleanings requested rather than as an absolute number.

Appendix F: Solution When $k_{o,H} - s_H(d_H^*) < k_{o,L} - s_L(d_L^*)$

Under this condition, the solution to $\max_y H$ is $y^* = \max\{0, \min(\gamma\alpha L, C - (1 - \gamma)\alpha L)\} = \min\{\alpha L, C\} - \min\{(1 - \gamma)\alpha L, C\}$. That is, the platform prioritizes low-value customers for assigning standby cleaners. The objective function has a similar structure to in the case $k_{o,H} - s_H(d_H^*) > k_{o,L} - s_L(d_L^*)$ in Appendix D. Conditioning on C , we start by simplifying the inner expectation on α .

$$\mathbb{E}_\alpha \left[\min\{(1 - \gamma)\alpha L, C\} (k_{o,L} - k_{o,H} - s_L(d_L^*) + s_H(d_H^*)) + \min\{\alpha L, C\} (k_{o,H} - s_H(d_H^*)) \right] + \\ (\gamma\mu_\alpha L(s_H(d_H^*) - k_H) + (1 - \gamma)\mu_\alpha L(s_L(d_L^*) - k_L) + (k_H\gamma L + k_L(1 - \gamma)L)),$$

where $\mu_\alpha = \mathbb{E}_\alpha[\alpha]$. We write the above as $\mathbb{E}_\alpha [\min\{(1 - \gamma)\alpha L, C\} M'_1 + \min\{\alpha L, C\} M'_2] + M_3$. The rest of the analysis follows from Appendix D.

Appendix G: Long-term Effect of No-shows on Customer Attrition

Complementing the one-period model in Section 5.2.1, we build a two-stage model demonstrating how no-shows can change platform size over time in this appendix.

We allow heterogeneity in customers' baseline attrition risk and denote customer type $\tau_r \in \{\mathcal{H}, \mathcal{L}\}$, where $\mathcal{H}$ and $\mathcal{L}$ represent the high attrition risk, and low attrition risk customers respectively. Suppose a fraction λ of customers have high attrition risk, while the remaining fraction $1 - \lambda$ have low attrition risk. Based on customer type, their additional attrition risks induced by no-shows are also different. In particular, when high-attrition-risk customers experience a no-show, they leave with probability 1; when high-attrition-risk customers do *not* experience a no-show, they leave with probability β_h . When low-attrition-risk customers experience a no-show, they leave with probability β_l ; when low-attrition-risk customers do *not* experience a no-show, they leave with probability 0. We assume homogeneous customer lifetime value, v , in the dynamic model for tractability.

The platform starts with L daily cleanings (customers) in the first period. For simplicity, we assume there is no market expansion. The no-show rate α is drawn from distribution F ; the same no-show rate applies in both stages. $V_1(C | L, \lambda, \beta_h, \beta_l)$ is the platform's present aggregate customer lifetime value when there are L initial customers in the first period, and the platform keeps C standby cleaners. The platform's expected present value is $\mathbb{E}_\alpha[V_1(C | L, \lambda, \beta_h, \beta_l)] = \lambda L \mathbb{E}_\alpha[V_{1,\mathcal{H}}(C | L, \lambda, \beta_h, \beta_l)] + (1 - \lambda) L \mathbb{E}_\alpha[V_{1,\mathcal{L}}(C | L, \lambda, \beta_h, \beta_l)]$, and V_{1,τ_r} are the values that the platform can collect from a customer of type τ_r in the first stage. Let β^0 be the discount factor. We write down the distributions of $V_{1,\mathcal{H}}$ and $V_{1,\mathcal{L}}$ as follows. Suppose when customer i was not impacted by no-show on the originally scheduled service date and time, the platform receives $k = p - m$. Particularly, "not impacted by no-show" happens in two scenarios: when the cleaning is *not* a no-show incident, or the cleaning is a no-show incident, but it is *fully* recovered by a remedy strategy: promptly restoring via a standby cleaner.

$$V_{1,i \in \tau_r}(C | L, \lambda, \beta_h, \beta_l) = \begin{cases} 0 & \text{if } i \text{ is impacted by no-show and left} \\ \beta^0 \mathbb{E}_\alpha[V_{2,i}(C | L', \lambda, \beta_h, \beta_l)] & \text{if } i \text{ is impacted by no-show and stayed} \\ k & \text{if } i \text{ is not impacted by no-show and left} \\ k + \beta^0 \mathbb{E}_\alpha[V_{2,i}(C | L', \lambda, \beta_h, \beta_l)] & \text{if } i \text{ is not impacted by no-show and stayed} \end{cases}$$

$p_{\mathcal{H},j}(L | C, \alpha)$, $p_{\mathcal{L},j}(L | C, \alpha)$ specify the probability of customer i being in each of the four states when the total number of cleanings is L . We suppress (C, α) below. In particular,

$$\begin{cases} p_{\mathcal{H},1}(L) = \alpha p_r(L), p_{\mathcal{L},1}(L) = \alpha p_r(L) \beta_l, & i \text{ is impacted by no-show and left;} \\ p_{\mathcal{H},2}(L) = 0, p_{\mathcal{L},2}(L) = \alpha p_r(L) (1 - \beta_l), & i \text{ is impacted by no-show and stayed;} \\ p_{\mathcal{H},3}(L) = (1 - \alpha p_r(L)) \beta_h, p_{\mathcal{L},3}(L) = 0, & i \text{ is not impacted by no-show and left;} \\ p_{\mathcal{H},4}(L) = (1 - \alpha p_r(L)) (1 - \beta_h), p_{\mathcal{L},4}(L) = 1 - \alpha p_r(L), & i \text{ is not impacted by no-show and stayed,} \end{cases}$$

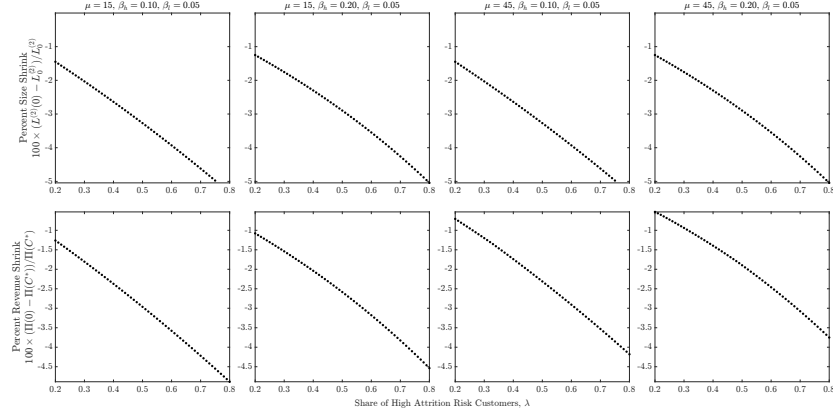


Figure G.1 Long-Run Effects of No-Shows on Platform Size and Platform Recovery

This figure illustrates the impact of no-shows on platform size and value in a two-stage model. The top panel shows how much the platform shrinks if no action is taken, as a function of the share of high-attribution-risk customers. Specifically, $L_0^{(2)}$ denotes platform size at the end of stage two with only baseline attrition, while $L^{(2)}(0)$ reflects size with both baseline and no-show-induced attrition, assuming the platform takes no action (i.e., retains no standby cleaners). The bottom panel shows the percent reduction in platform value relative to the policy that optimally retains standby cleaners in both stages. $\Pi(0)$ represents the platform value under no-shows without intervention, while $\Pi(C^*)$ is the value with optimal standby capacity. In all subplots, aside from μ , β_h , and β_l (varied as shown in titles), parameters include $\beta^0 = 0.9$, $v = 300$, $\alpha \sim \text{Unif}$, and $(L, p, m, \mu_\alpha, \sigma_\alpha)$ take values as in Figure 1.

where $p_r(L) := p_r(L|C, \alpha)$ denotes the probability that customer i is *not* covered by a standby cleaner when a no-show occurs, in particular, $p_r(L|C, \alpha) = \max\{0, 1 - C/(\alpha L)\}$.

The number of customers changes in period 2, as some customers may have already left the platform due to the service failure they experienced in period 1. Among high-attribution-risk customers, the probability of leaving after the first period is $p_{\mathcal{H},a}(L) = \alpha p_r(L) + (1 - \alpha p_r(L))\beta_h$ while for low-attribution-risk customers, the probability is $p_{\mathcal{L},a}(L) = \alpha p_r(L)\beta_l$. The number of customers (cleanings) in the second period is then the sum of two binomial random variables, $L' \sim \text{Bin}(\lambda L, 1 - p_{\mathcal{H},a}(L)) + \text{Bin}((1 - \lambda)L, 1 - p_{\mathcal{L},a}(L))$.

$$V_{2,i \in \tau_r}(C|L', \lambda, \beta_h, \beta_l) = \begin{cases} 0 & \text{with probability } p_{\tau_r,1}(L') \text{ when } i \text{ is impacted by no-show in stage 2 and left} \\ v & \text{with probability } p_{\tau_r,2}(L') \text{ when } i \text{ is impacted by no-show in stage 2 but stayed} \\ k & \text{with probability } p_{\tau_r,3}(L') \text{ when } i \text{ is not impacted by no-show in stage 2 and left} \\ k+v & \text{with probability } p_{\tau_r,4}(L') \text{ when } i \text{ is not impacted by no-show in stage 2 and stayed} \end{cases}$$

Therefore, $\mathbb{E}_\alpha[V_{2,\tau_r}(C|L', \lambda, \beta_h, \beta_l)] = v p_{\tau_r,2}(L') + k p_{\tau_r,3}(L') + (k+v) p_{\tau_r,4}(L')$.

Denote by μ the present-value cost per standby cleaner over the two-stage horizon. The platform chooses the number of standby cleaners C to maximize its present value: $\max_{C \in \mathbb{Z}^+ \cup \{0\}} \mathbb{E}_\alpha[V_1(C|L, \lambda, \beta_h, \beta_l)] - \mu C$. We solve for C and conduct numerical analysis on two metrics: how platform size evolves over time due to no-shows when no action is taken, and how much value can be recovered through remedy strategies such as retaining standby cleaners. The results are shown in Figure G.1, assuming an average no-show rate of 3.5% per period. Over two periods, this leads to up to 5% reduction in platform size if unaddressed. Without an optimal remedy strategy, the platform incurs significant revenue losses. Remedy strategies such as retaining standby cleaners can recoup a significant share of such losses.

Appendix Endnotes

¹To construct the final analysis sample, we apply several filters. The main exclusions are cities with limited business (1,026 cleanings), cleanings performed by teams with more than one cleaner (1,896 cleanings), and orders placed by corporate customers (6,344 cleanings).

²Our findings remain robust when these observations are included.

³We thank the anonymous reviewer for highlighting this potential confounding factor.

⁴For this analysis, we excluded 179 no-show cleanings for which the no-show cleaner ID is missing.

⁵Our sample includes 1,045 no-show cleanings that were rescheduled in the same day.

⁶For no-show events, the originally matched cleaner ID is missing for 179 cleanings; these cleanings are excluded from our predictive model. For one city in our sample, the platform did not collect cleaner district data. The performance of our model remains consistent when the distance metric between the customer and the cleaner is excluded from the prediction model.

⁷We select our main specification by choosing among eight candidate covariate sets based on BIC following Bachmann et al. (2021). See Appendix B.1 for details on covariate definitions and candidate covariate sets.

⁸Our counterfactual analysis is robust when conducted by customer cohort. To ensure better comparability across customers, we define cohorts by quarter following Fader and Hardie (2001): customers whose first transaction occurred during the first quarter (December 2018 to March 2019) form the first cohort, and customers whose first transaction occurred during the second quarter (April 2019 to June 2019) form the second cohort. In cohort-specific analyses, eliminating no-shows for half a year would increase CLV in both cohorts.

⁹The platform pays 75% of the total cost per cleaning. Total cost is price paid by customers plus the total discount offered.

¹⁰In particular, for each customer, we sum their spending in the first 52 weeks, and add the predicted CLV over the subsequent 8 holdout weeks to form a CLV distribution.

References

- Austin PC (2009) Balance diagnostics for comparing the distribution of baseline covariates between treatment groups in propensity-score matched samples. *Statistics in Medicine* 28(25):3083–3107.
- Bachmann P, Kuebler N, Meierer M, Naef J, Oblander E, Schilter P (2022) *CLVTools: Tools for Customer Lifetime Value Estimation*. URL <https://github.com/bachmannpatrick/CLVTools>, r package version 0.9.0.
- Bachmann P, Meierer M, Näf J (2021) The role of time-varying contextual factors in latent attrition models for customer base analysis. *Marketing Science* 40(4):783–809.
- Boyd SP, Vandenberghe L (2004) *Convex Optimization* (Cambridge University Press).
- Cohen MC, Fiszer MD, Kim BJ (2021) Frustration-based promotions: Field experiments in ride-sharing. *Management Science* 68(4):2432–2464.
- Colombo R, Jiang W (1999) A stochastic rfm model. *Journal of Interactive Marketing* 13(3):2–12.
- Conley TG, Hansen CB, Rossi PE (2012) Plausibly exogenous. *Review of Economics and Statistics* 94(1):260–272.
- Fader PS, Hardie BG (2001) Forecasting repeat sales at cldnow: A case study. *Interfaces* 31(3_supplement):S94–S107.
- Fader PS, Hardie BG, Lee KL (2005) Rfm and clv: Using iso-value curves for customer base analysis. *Journal of Marketing Research* 42(4):415–430.
- Frake J, Harmon D (2023) Intergenerational transmission of organizational misconduct: Evidence from the chicago police department. *Management Science* 70(6):3856–3878.

- Guo Y, Zhang Y, Goh KY, Peng X (2025) Can social technologies drive purchases in e-commerce live streaming? an empirical study of broadcasters' cognitive and affective social call-to-actions. *Production and Operations Management* 34(12):4039–4059.
- Jain N, Tan TF (2022) M-commerce, sales concentration, and inventory management. *Manufacturing & Service Operations Management* 24(4):2256–2273.
- Jeni LA, Cohn JF, De La Torre F (2013) Facing imbalanced data—recommendations for the use of performance metrics. *2013 Humaine Association Conference on Affective Computing and Intelligent Interaction (ACII)*, 245–251 (IEEE).
- Kalbfleisch JD, Prentice RL (2002) *The Statistical Analysis of Failure Time Data* (John Wiley & Sons).
- Lemmens A, Gupta S (2020) Managing churn to maximize profits. *Marketing Science* 39(5):956–973.
- Marinoni A, Roche MP (2025) You've got mail! the late 19th century u.s. postal service expansion, firm creation, and firm performance. *Management Science* 71(9):7223–7243.
- Milgrom P, Shannon C (1994) Monotone comparative statics. *Econometrica* 62(1):157–180.
- Neslin SA, Gupta S, Kamakura W, Lu J, Mason CH (2006) Defection detection: Measuring and understanding the predictive accuracy of customer churn models. *Journal of Marketing Research* 43(2):204–211.
- Ninh A, Shen ZJM, Lariviere MA (2020) Concavity and unimodality of expected revenue under discrete willingness to pay distributions. *Production and Operations Management* 29(3):788–796.
- Oblander ES, McCarthy D (2021) How has covid-19 impacted customer relationship dynamics at restaurant food delivery businesses. *Mark. Sci. Inst. Work. Pap. Ser* 23(2139):35.
- Schmittlein DC, Morrison DG, Colombo R (1987) Counting your customers: Who-are they and what will they do next? *Management Science* 33(1):1–24.
- Shaked M, Shanthikumar JG (2007) *Stochastic Orders* (Springer).
- Song H, Tucker AL, Murrell KL, Vinson DR (2017) Closing the productivity gap: Improving worker productivity through public relative performance feedback and validation of best practices. *Management Science* 64(6):2628–2649.
- Staats BR, Dai H, Hofmann D, Milkman KL (2016) Motivating process compliance through individual electronic monitoring: An empirical examination of hand hygiene in healthcare. *Management Science* 63(5):1563–1585.
- Xu Y, Lu B, Ghose A, Dai H, Zhou W (2023) The interplay of earnings, ratings, and penalties on sharing platforms: An empirical investigation. *Management Science* 69(10):6128–6146.