



GROWTH THAT COMPOUNDS

One customer model cannot predict every outcome

Next purchase, partial defection, and profitability are different outcomes: choose the model and validation question before trusting one customer score.

Sinan Isoglu, MBA (Quantic)

Commercial growth leader, lecturer and doctoral researcher

PUBLISHED

1 September 2026

READING TIME

7 min

WORDS

1,483

<https://isoglu.com/insights/journal/one-customer-model-cannot-predict-every-outcome/>

Management summary	2
Why does a unified customer value label conceal three distinct behavioral outcomes?	3
What does random forest benchmarking prove regarding multi-outcome customer predictions?	4
Why does statistical predictive importance diverge from causal operational leverage?	5
What disciplined review sequence evaluates churn, expansion, and margin separately?	5
Which empirical transfer test must precede deploying a predictive model across cohorts?	5
Which three predictive customer modeling assumptions must data leaders reject?	6
Where are the methodological boundaries of enterprise customer modeling?	6
References	7

MANAGEMENT SUMMARY

A customer model can predict a next purchase without explaining partial defection or profitability evolution. Lariviere and Van den Poel study 100,000 customers in a financial-services setting and model those three outcomes with random forests and regression forests. In their reported applications, the forest techniques fit and validate better than ordinary linear and logistic benchmarks, while the importance of the same variables differs across buying, defection, and profitability. This article turns that result into a synthetic multi-outcome model card. It separates predictive fit, variable importance, validation, and intervention value, and does not present the study as a universal CRM benchmark or a causal retention result.

Keywords: Customer outcome prediction · Random forest · Customer profitability · Partial defection

A customer model that predicts a next purchase has not necessarily explained defection or profitability.

The short answer is that next purchase, partial defection, and profitability evolution are different outcomes, so they need separate outcome definitions, validation questions, and decision uses. One score can be convenient. It can also hide the fact that a variable that helps predict buying is not the variable that matters most for defection or profit.

Lariviere and Van den Poel study a real-life sample of 100,000 customers from a large European financial- services company. They use random forests for binary outcomes and regression forests for outcomes with a linear dependent variable. In the reported applications, both forest techniques fit and validate better than ordinary linear and logistic regression benchmarks. The importance of the same explanatory variables also differs across the three outcomes.

Malthouse and Blattberg (2005) add a separate Customer Lifetime Value prediction problem, including the cost of misclassifying future high-value customers. Lemmens and Gupta (2020) add an intervention problem: retention targeting should rank incremental profit rather than churn risk or response alone. Prediction and intervention can share data, but they are not the same outcome or decision.

Why does a unified customer value label conceal three distinct behavioral outcomes?

“Will this customer buy?” is not the same as “Will this customer partially defect?” or “How will this customer’s profitability evolve?” The questions can be related, but a related outcome is not a substitute for the declared target.

Table 1 Why does a unified customer value label conceal three distinct behavioral outcomes?

Outcome	Decision question	What can go wrong if it is used as a proxy
Next purchase	Which customers are likely to buy again within the horizon?	A purchase can be small, unprofitable, or unrelated to retention
Partial defection	Which customers may reduce their relationship or activity?	A reduced activity pattern can be missed by a binary churn label
Profitability evolution	Which customers’ profit contribution may change?	Revenue or purchase count can omit cost and margin

Table from this essay. Sources and interpretation are given in the article.

The first discipline is to name the event. The second is to name the time horizon. The third is to say what action the prediction will support. A model can be accurate for one event and irrelevant to another because it is answering a different question.

Table 1 The multi-outcome customer model card

Synthetic outcome	Model family	Validation question	Influential variable class	Decision use	Transfer risk
Next purchase	Classification forest	Does the model separate future buyers from non-buyers?	Past purchase behavior	Prioritize a follow-up test	Purchase value is not profit
Partial defection	Classification forest	Does the model detect a declared reduction in activity?	Intermediary or channel behavior	Review relationship change	Defection boundary differs by channel
Profitability evolution	Regression forest	Does the model predict change in the declared profit measure?	Past behavior and cost boundary	Review economic exposure	Revenue-only validation misleads

Author's synthetic model-review framework grounded in Lariviere and Van den Poel (2005), Malthouse and Blattberg (2005), and Lemmens and Gupta (2020). The rows are illustrative and do not describe a live customer model.

The table is a review device. It does not report performance. It makes an outcome substitution visible: purchase behavior may help identify repeat buying, but it does not automatically measure profitability. Intermediary behavior may matter for defection, but that importance belongs to the source's setting and model.

What does random forest benchmarking prove regarding multi-outcome customer predictions?

Lariviere and Van den Poel use random forests for binary classification and regression forests for models with a linear dependent variable. They compare the techniques with ordinary linear and logistic regression benchmarks in estimation and validation samples. The reported result is better fit for both forest techniques in the applications studied.

That is a model comparison, not a universal algorithm ranking. Malthouse and Blattberg (2005) add a forecast-horizon and misclassification boundary to the Customer Lifetime Value use case. Better fit in one financial-services application does not tell a different organization which model to use. The outcome definition, data quality, time horizon, missingness, class balance, cost boundary, and decision use can all change the comparison.

The source also reports that the same explanatory variables have different impacts on buying, defection, and profitability. Past customer behavior is more important for repeat purchasing and favorable profitability evolution in the reported evidence, while the intermediary's role has a greater impact on defection proneness. The point is not to memorize those importance statements as a portable feature ranking. The point is to ask whether the model is allowed to use one outcome's drivers to stand in for another.

Why does statistical predictive importance diverge from causal operational leverage?

Variable importance answers a predictive question inside a model and sample. It does not show that changing the variable will create the predicted outcome. A past purchase pattern can be informative without being a lever a manager can change. An intermediary variable can identify a defection pattern without proving that changing the intermediary relationship will prevent it.

This distinction protects the boundary between prediction and intervention. Lemmens and Gupta (2020) show that retention targeting adds an incremental-profit question that predictive fit alone cannot answer. A model can help identify which customers deserve a closer review. A retention campaign needs an intervention, a comparison condition, a cost boundary, and an outcome horizon. Predictive fit alone does not provide the counterfactual.

The distinction also applies to profitability. A model can predict profitability evolution while using variables that encode historical cost or service differences. If the decision asks where to invest, the model review must show whether those inputs are available before the investment and whether the resulting action changes the economic outcome.

What disciplined review sequence evaluates churn, expansion, and margin separately?

Before a customer model is allowed to steer a decision, record:

- 1 the exact outcome and event definition;
- 2 the prediction horizon and observation window;
- 3 the model family and benchmark model;
- 4 the estimation and validation samples;
- 5 the measure of fit and the cost of misclassification;
- 6 the variables that matter for this outcome;
- 7 the variables that are merely correlated signals;
- 8 the action the score is permitted to trigger;
- 9 the evidence that an action changes the outcome rather than only predicts it.

The eighth field is easy to omit. A score can be useful for triage, explanation, service planning, or intervention selection. Those are not the same decision. A model that predicts partial defection may trigger a relationship review. It does not automatically authorize a discount, a service change, or a retention offer.

Which empirical transfer test must precede deploying a predictive model across cohorts?

If a model or finding is transferred to another customer setting, ask five questions:

Table 3 Which empirical transfer test must precede deploying a predictive model across cohorts?

Transfer question	Why it matters
Is the outcome defined the same way?	“Defection” and “profitability” can have different boundaries
Is the observation horizon comparable?	Short-term buying and long-term value are not interchangeable
Are the input variables available before the action?	A post-outcome signal cannot support a pre-outcome intervention
Is the cost boundary comparable?	Revenue, margin, and profit can reverse a priority
Is the decision use the same?	Prediction for triage is not prediction for treatment

Table from this essay. Sources and interpretation are given in the article.

If the answers are unknown, the source can still motivate a test. It cannot support a portable benchmark.

Which three predictive customer modeling assumptions must data leaders reject?

First, do not say “the model predicts customer value” when the target is only next purchase. Name the outcome.

Second, do not say “random forests are better” without naming the comparison, sample, outcome, and validation boundary. The source reports applications, not a universal winner.

Third, do not say “the important variable should be changed” merely because it predicts the result. Predictive importance is not intervention evidence.

For adjacent decisions, compare the customer-lifetime-value forecast with the renewal evidence boundary.

Where are the methodological boundaries of enterprise customer modeling?

Lariviere and Van den Poel provide a 100,000-customer financial-services application in which random and regression forests fit and validate better than ordinary benchmarks, while variable importance differs across next purchase, partial defection, and profitability evolution. Malthouse and Blattberg provide a forecast-horizon and misclassification boundary, while Lemmens and Gupta provide an intervention-profit boundary. The three studies do not provide a universal CRM benchmark or a causal retention result. The synthetic model card is an author-owned translation. Its job is to keep outcome, model, validation, predictor interpretation, and decision use from collapsing into one customer score.

REFERENCES

- Lariviere, B., & Van den Poel, D. (2005). Predicting customer retention and profitability by using random forests and regression forests techniques. *Expert Systems with Applications*, 29(2), 398-405. DOI
- Malthouse, E. C., & Blattberg, R. C. (2005). Can we predict customer lifetime value? *Journal of Interactive Marketing*, 19(1), 2-16. DOI
- Lemmens, A., & Gupta, S. (2020). Managing churn to maximize profits. *Marketing Science*, 39(5), 956-973. DOI

END OF ESSAY

HOW TO CITE THIS ESSAY

Isoglu, S. (2026, September 1). One customer model cannot predict every outcome. Sinan Isoglu.
<https://isoglu.com/insights/journal/one-customer-model-cannot-predict-every-outcome/>

KEEP READING

Growth that compounds

The 20-55 rule: customer prioritization misclassifies the portfolio

<https://isoglu.com/insights/journal/the-20-55-rule-customer-prioritization-misclassifies-the-portfolio/>

Growth that compounds

A customer P&L needs a cost boundary

<https://isoglu.com/insights/journal/a-customer-p-and-l-needs-a-cost-boundary/>

Growth that compounds

Long customer relationships can be low-profit

<https://isoglu.com/insights/journal/long-customer-relationships-can-be-low-profit/>



Sinan Isoglu, MBA (Quantic)

Commercial growth leader, lecturer and doctoral researcher