

Is There a Gold Standard or a Need for a City-Centric Approach to Sales Tax Revenue Forecasting

Sarah E. Larson^{i,c}, Michael R. Overtonⁱⁱ

The accuracy of sales tax revenue forecasting is essential for local governments, as they rely on these forecasts to develop their annual budgets. Previous research has focused on identifying gold-standard forecasting methods with high average accuracy across multiple cities. However, such approaches may still produce inaccurate predictions for specific municipalities, making this scholarship less relevant to practitioners. Our research addresses the gap in the existing literature by focusing on the relative accuracy of forecasts from the municipal perspective rather than the overall average accuracy across all municipalities—a city-centric approach—to identify variations in various machine learning and traditional revenue forecasting methods. Here, we show that following the steps of PREE: (P) prepare, (R) run multiple models, (E) evaluate against benchmarks, and (E) evaluate overall performance can help to maximize the accuracy of sales tax revenue forecasting at the municipal level. The high variability in model performance across municipalities highlights the risks associated with relying on a single gold-standard forecasting approach. Instead, practitioners should focus on forecasting processes, such as PREE.

Keywords: Forecasting Accuracy, Machine Learning, Revenue Forecasting, Sales Taxation

Local governments rely on accurate revenue forecasts to develop budgets, develop long-term plans, and monitor their fiscal health. Ideally, forecasts impose budgetary discipline by acting as resource constraints while governments construct annual operating budgets. While public budgeting and finance forecasting scholarship has focused on forecasting method accuracy and politically induced forecasting bias, no scholarship has looked at the risk associated with effective forecasting methods from the practitioner's perspective. Accuracy in revenue forecasting has focused on the overall average accuracy of various forecasting techniques. Our research addresses this critical gap by evaluating sales tax forecasting methods at the individual

ⁱ Department of Political Science, Miami University.  <https://orcid.org/0000-0002-9644-2019>.

ⁱⁱ Department of Politics and Philosophy, University of Idaho.  <https://orcid.org/0000-0002-1800-9164>.

^c Corresponding Author: larsonse@miamioh.edu.

city level—aligning our study with the practical application of city forecasting practices. We examine how forecasting methods demonstrating high average accuracy across multiple cities may still produce inferior predictions for specific municipalities. This city-centric approach aims to enhance local government’s ability to choose forecasting methods tailored to their unique characteristics through a robust forecasting procedure called PREE: (P)repare, (R)un multiple models, (E)valuate against benchmarks, and (E)valuate overall performance, thereby minimizing the risk of using relatively inaccurate projections and improving overall fiscal planning efficacy.

Most sales tax forecasting studies have explored the accuracy of revenue forecasting methods by examining the political and institutional factors that introduce bias into the forecasting process or by comparing forecasting models across multiple units, such as cities seeking the best model or a gold standard of forecasting. This study focuses on the latter, where a gold standard will produce revenue forecasted figures with the greatest average accuracy. While revenue forecasting accuracy is never going to be perfect, as revenue forecasting is part art and part science, a gold standard would streamline the time and effort spent in developing forecasting recommendations. Any municipality could collect and prepare sales tax data, run the current gold standard revenue forecast, and be guaranteed maximum forecasting accuracy—a reliable and evidence-based set of processes to streamline revenue forecasting for state and local governments.

While it is understandable, searching for a gold standard in forecasting poses challenges. Larson and Overton (2024) suggested that focusing simply on revenue forecasting methods was a short-sighted approach to increasing forecasting accuracy. Instead, they suggested that state and local officials who wish to forecast sales tax revenue should focus on preparing or preprocessing the data before applying the revenue forecasting technique, as selecting appropriate preprocessing steps improves forecasting accuracy considerably more than any model.

In addition, the existing revenue forecasting literature has limited practical application for city-level practitioners because pursuing a gold-standard forecasting model requires unrealistic assumptions and is evaluated using scholarly rather than practitioner-oriented framings. First, the gold standard approach is typically evaluated by averaging the accuracy of forecasts across multiple units such as cities, counties, or other government entities. The model that generates the most accurate forecasts *on average* becomes the new gold standard and is recommended as a best practice for all cities. This process would work if the underlying conditions that drove sales tax revenue were the same. It should be noted that there is no explicit public administration scholarship on pursuing a “gold standard.” However, the term is a useful proxy for scholarship that focuses on finding the most accurate forecast on average. However, sales tax revenue is generated by highly localized factors (Makridakis et al., 2020), such as spending patterns, population size and growth, wages, weather, and tax rates. What captures the sales tax revenue dynamics for one city is very unlikely to capture it in other cities. Second, outlying forecasts can skew a model’s average accuracy, biasing the average forecast. One or more extremely accurate or inaccurate forecasts can result in erroneous recommendations on a city-by-city basis.

Instead, a *city-centric* approach is required to help local practitioners develop forecasting practices that evaluate and benchmark revenue forecasts to identify the most locally accurate forecast rather than simply implementing the latest gold standard method. By adopting a city-centric approach, we aimed to reframe the analysis of forecasting accuracy in a practitioner-friendly manner. Like other studies, we ran a variety of models across multiple cities. However, before averaging the results, we ranked their relative accuracy for each city in our sample. The

results captured the noisy reality of local sales tax forecasting. Therefore, this study addresses the following research question: How does a city-centric approach to sales tax forecasting impact the evaluation of the relative accuracy of various preprocessing and revenue forecasting techniques for municipalities attempting to forecast sales tax revenue?

To address this research question, we used sales tax data from roughly 1,000 unique cities in Texas. Since we wanted to generate city-centric insight, we first outlined PREE as a forecasting process for practitioners. Using this framework, we compared the performance of various data preprocessing methods (i.e., preparation) and forecasting methods (i.e., running multiple models) against benchmarks (i.e., evaluating against benchmarks) and all model–preprocessing (MP) combinations for each city (i.e., evaluating overall performance). PREE equips practitioners with a set of practical steps to follow systematically each time revenue forecasting is undertaken by the municipality to increase the odds of accurate forecasts.

To overcome the problems associated with the average accuracy approach that can lead to inaccurate city-level forecasts, we used the relative accuracy of each forecast by ranking each forecast's accuracy for every city in our sample. This approach enabled us to evaluate our forecasts from a city-centric perspective rather than determining when generally effective forecasts became relatively inaccurate. Understanding the odds of the best practice approach failing a municipality and other common preprocessing and revenue forecasting techniques helped to ground this manuscript in the noisy reality of local sales tax revenue forecasting. Through practicing PREE and highlighting the presence of deviations, we emphasize the importance of a city-centric approach in future revenue forecasting literature, particularly for practical use by practitioners.

The remainder of the paper is organized as follows: First, the existing literature on forecasting accuracy and accuracy volatility is presented. Second, the data and sources are discussed. Third, a summary of the city-centric framework for local government revenue forecasting is presented. Next, a summary of the methods employed in this research is presented. Subsequently, the results are presented. Finally, a conclusion and a broader overview of the more significant impacts of the findings are presented.

Literature Review

Accuracy and its importance in revenue forecasting have been central to scholars' discussions within the literature over the last 40 years (Bretschneider et al., 1989; Fullerton, 1989; M.M. Rubin et al., 2019). The early literature focused on the political nature of the forecasting process, highlighting the impacts of individuals and the more extensive political process of state and local governments on budgeting, which in turn affects forecasting accuracy (Bretschneider & Gorr, 1992).

Subsequent literature focused on the impacts of recessions on revenue forecasting accuracy. The recession-protecting desire to acquire rainy-day funds, acting against political forces, suggests that hiding the presence of the funds or the need to spend them created a struggle for forecasting accuracy. The political impact of citizens knowing of any surplus acted against the desire to retain funds in times of economic success for use in times of recession (Rodgers & Joyce, 1996). Revenue forecasting during times of economic recession can be challenging for even the most seasoned forecaster (Mikesell, 2018).

Some lines of literature have suggested that a lack of forecasting accuracy was due to political factors and individual forecaster biases (Williams, 2012). Political factors that influence the accuracy of state revenue forecasting include budget timing within the electoral cycle, the presence of incumbents, and the political party in control (Brogan, 2012). The underlying impact of individual politicians on revenue forecasting accuracy stems from three characteristics: politicians' aversion to risk, a desire to have a perception that finances are managed properly, and flexibility in budgetary authority to increase spending if necessary to obtain votes for reelection (Tversky & Kahneman, 1992; Krause, 2012; Rodgers & Joyce, 1996).

However, findings on the impact of political factors on forecasting accuracy were mixed. Mocan and Azad (1995) found no impact of the dominant political party upon forecasting accuracy. At the local level in Florida, finance officers were not receiving the necessary political and bureaucratic scrutiny to enhance forecasting accuracy (Frank & Zhao, 2009).

An internal focus of the research during the period also explored the importance of actions such as the use of independent forecasting agencies, budget preparation, and internal accounting reporting on the accuracy of forecasting. This research explored the impact of preparation and internal accounting reporting on forecasting accuracy. Cassar and Gibson (2008) found a strong relationship between internal accounting report preparation and forecasting accuracy. However, they found a less statistically significant relationship between budget preparation and forecasting accuracy. Establishing an independent forecasting agency and technical workgroups improved forecasting accuracy in Washington State (Deschamps, 2004). Lorenz and Homburg (2018) found that analysts with weak forecasting abilities often stopped forecasting revenues, as accuracy may be tied to promotion, while inaccuracy may be tied to termination.

During this period, cross-country differences in revenue forecasting accuracy were also explored. Buettner and Kauder (2010) found cross-country differences based on the relative importance of various taxes within the overall revenue diversification of particular countries' tax portfolios, such as the relative importance of corporate or personal income taxation within the countries' revenue structures. Buettner and Kauder also found that the timing of forecasting was important in explaining accuracy. In other cases, Mikesell and Ross (2014) argued that a focus on forecasting accuracy was inappropriate; rather, the focus should be on the political nature of obtaining consensus on the revenue forecasting figures during the budgeting process.

During a subsequent refocusing, the revenue forecasting literature shifted to the accuracy of various methods, particularly on which methods allow for the greatest forecasting accuracy of future revenue streams. The early focus during this shift was on the use of computer technology, the experience of individual forecasters, and specific methods (Kong, 2007; MacManus & Grothe, 1989). Kong (2007) highlighted the need for more preparation of revenue officials in California counties to use sophisticated techniques. A lack of exposure to these sophisticated techniques was attributed to the training of local government budget officers in public administration programs in the United States (Reddick, 2004).

The shift to a focus on the accuracy of various methods was driven by several streams of research that identified a direct relationship between revenue forecasting inaccuracy and municipal fiscal stress (Caiden & Wildavsky, 1980; Chapman, 1982; Forrester, 1991; I. S. Rubin, 1987). Municipalities experiencing fiscal stress were more likely to rely on more sophisticated forecasting techniques (MacManus & Grothe, 1989); the switch to these techniques increased future revenue forecasting accuracy. Scholars disagreed on the direct relationship between revenue forecasting accuracy or inaccuracy and fiscal stress. Still, the literature has

broadly focused on the accuracy of various forecasting methods used by state and local government officials. The focus on accuracy is crucial for revenue forecasting, which helps streamline the work of local government officials.

As technology has become increasingly accessible to local government officials, machine-learning techniques have emerged as methods for revenue forecasting (Buxton et al., 2019). While the K-nearest neighbor algorithm has been in the scholarly literature since the 1960s (Clover & Hart, 1967), the computational requirements to use these algorithms prevented broad-scale municipal implementation at that time. Scholars have grappled with using machine learning approaches (Carmody & Wiipongwii, 2018; Hansen & Nelson, 1997, 2002; Muh & Jang, 2019; Voorhees, 2006) rather than traditional (time-series or causal-like approaches) (Williams & Calabrese, 2016) to increase forecasting accuracy with mixed results (for example, Buxton et al., 2019; Chung et al., 2022). The accuracy of traditional versus machine-learning techniques was also compared over various forecasting periods (monthly versus quarterly) to determine whether the most accurate approach changed when the forecasting period changed (Williams & Kavanagh, 2016).

This focus came in part from a desire to streamline the practice of forecasting revenues for state and local governments—a desire to suggest that the gold standard method represents the best approach for forecasting a specific type of revenue under certain conditions. If the gold standard can be identified, revenue analysts can apply these techniques to their municipalities' or states' data and have confidence in the most accurate results.

Accuracy in revenue forecasting and the desire for accuracy are driven by both internal and external factors (Reitano, 2018). External factors, such as recessions and the political nature of budgeting, drive accuracy, as well as the prominent role of forecasting within the larger budget. Internal factors, such as sophisticated machine learning methods or even the forecaster's experience, impact revenue forecasting accuracy.

However, public forecasting scholarship has not adequately addressed the needs of practitioners who seek to achieve forecasting accuracy using highly localized time-series data with unique properties and underlying data-generating processes. The disconnect between practitioner needs and public revenue forecasting scholarship may stem from the fact that individual-level analysis is not always aligned with the demands for publication in academic journals and the academic expectations for promotion and tenure. Therefore, the public forecasting scholarship comprises many large-N studies that advise the average municipality. Advice for the average municipality can lead to inaccurate city-level forecasts, as not every municipality is representative of the average. A city-centric approach is required to reframe how public sector forecasting research is conducted and presented to practitioners.

Benchmarking in Revenue Forecasting

Benchmarking in the tax revenue forecasting literature is not widely discussed. In forecasting, simple methods such as naïve, seasonal naïve, mean, and drift models can be calculated and serve as benchmarks to compare the accuracy of different, more sophisticated approaches (Hyndman & Athanasopoulos, 2018). These benchmarking methods are easy to calculate and make few assumptions about the data. Therefore, practitioners prefer benchmarking forecasts due to their ease of use and minimal assumptions (M. M. Rubin et al., 1999).

When evaluating the accuracy of a forecasting method, benchmarking methods offer an alternative and practical means of assessment. The benchmarking methods' mean absolute

percent error (MAPE) measure scores provide a lower bound with which various forecasts and scenarios can be compared. Exceeding the accuracy of the benchmark models is a vital step in selecting a revenue forecasting model. Furthermore, comparing methods to benchmark them prevents analysts from unnecessarily using a complex model when a simple forecast would have provided equally accurate predictions. If the purpose of the budget is to make internal management decisions (Forrester, 1991), expediency in predictive accuracy is essential. In contrast, if new modeling strategies are being considered, benchmarks provide a valuable and easy-to-implement check in the process.

Data Discussion and Why Texas?

Building on the work of Larson and Overton (2024), our study focuses on cities in Texas. The dataset encompasses over 1,000 cities tracked across 16 years, yielding a substantial sample for evaluating the city-centric accuracy of forecasting methodologies. Texas presents a particularly advantageous sample frame due to the stability and uniformity of its municipal sales tax rates over time. By restricting our analysis to Texas, we also ensure consistency in institutional frameworks governing collection practices, tax base definitions, and other administrative policies across all cross-sectional units under examination. In addition, cities within Texas vary dramatically in the amount of sales tax revenues collected and municipality size, with Houston having over 2.3 million residents and other cities having fewer than 1,000 residents (Texas State Comptroller, 2022.).

Rate changes would increase volatility and make it difficult to compare the accuracy of forecasts across time. Therefore, we needed a sample with relatively uniform rates over a long period. Texas imposes an 8.25% ceiling on the total sales tax, with 6.25% earmarked for the state's general fund (Texas State Comptroller, 2022). Of the remaining 2%, 1% is authorized for county and city use. The remaining 1% sales tax can be levied for special districts, such as economic development corporations, public transit authorities, and emergency services districts. The state initially collects sales tax revenue on a monthly, quarterly, or annual basis, contingent upon businesses' operational scales, with the frequency of filing requirements increasing as businesses collect more monthly sales tax revenue. Later, each lower-level unit of government levying a sales tax is allocated its portion of the total revenue.

A City-Centric Approach to Forecasting—PREE

As highlighted by the prior section, there is a lack of literature focusing on revenue forecasting at the individual city level. Scholarship analyzing the average forecasting accuracy or inaccuracy of various methods provides municipal forecasters with risk recommendations. Best practices offer little utility if they do not apply to municipalities when *forecasting their revenue*. The underlying forces that drive sales tax revenue are volatile and highly localized, making a one-size-fits-all approach convenient for scholars but a poor guide for practitioners. Therefore, this research presents a series of practical steps forecasters can use to evaluate and select the most locally accurate forecasting model.

These practical steps can be remembered using the acronym PREE, which aligns with the approach of testing by sampling in forecasting. Practicing PREE is not a one-time event but

rather a set of practical steps to follow systematically each time revenue forecasting is undertaken by the municipality to increase the odds of accurate forecasts.

First, the P in PREE stands for preparing your data for analysis. Larson and Overton (2024) demonstrated that the preparation steps of logging time-series data and performing inverse hyperbolic sine (IHS) transformations resulted in greater accuracy in sales tax revenue forecasting. It is crucial in this phase to hold out the latest 2 years of data to test the accuracy of the developed models (Williams & Kavanagh 2016). Second, the R in PREE represents the second step of running multiple types of revenue forecasts. As there is no single best way to forecast revenues for an individual municipality over time consistently, state and local revenue forecasters must run multiple revenue forecasting approaches. Third, the first E in PREE represents the need to evaluate the sales tax revenue forecasting results against benchmarks. Finally, the second E in PREE stands for the need to evaluate the overall forecasting accuracy of the revenue forecasts produced during the R and the first E stages of PREE.

Suppose state and local officials wishing to forecast sales tax revenues followed PREE. In that case, they would be following best practices in revenue forecasting by focusing not only on the forecasting methodology but also on a more extensive, holistic, and city-centric approach towards revenue forecasting. In the following sections, we frame our Methods and Results sections, which are organized using the PREE framework.

PREE: Methods

Prepare

Preparation involves identifying data, preprocessing it, and holding out a portion of the data for evaluative purposes.

Identification: Monthly sales tax collection data for every city in Texas were collected between January 1991 and December 2017 using the Texas Comptroller's website. In 2020, the Texas State Comptroller changed Comptroller Rule § 3.334, which governs sourcing rules for internet sales by in-state retailers in Texas. These sourcing changes sparked legal challenges to the rule change in the *City of Round Rock v. Hegar*, with additional changes made to Comptroller Rule § 3.334 in January 2023. We limited our data collection to before the changes to avoid including these rule changes within our data.

Unfortunately, continuous monthly data were not available for every city in Texas for the entire timeframe. Only cities with complete time series were included in the analyses, resulting in 822 cities with complete monthly data, 976 cities with complete quarterly data, and 1,005 cities with complete yearly data. While it was rare, some cities in Texas changed their sales tax rates during the period under study. Cities that did not have uniform sales tax rates throughout the study were weighted to ensure that the rate changes would not affect the forecasting accuracy.

Preprocessing: All data were inflation-adjusted to 2017 to ensure purchasing power comparability over time, and IHS was transformed to produce normalized distributions, as recommended by Larson and Overton (2024). Three preprocessing variations were generated for the monthly and quarterly forecasts: (a) no additional preprocessing steps, (b) seasonally adjusted data, and (c) detrended data. Seasonal adjustments and detrended preprocessing steps were calculated using multiplicative classical decomposition.

Table 1. Model Descriptions

Model	Full Name	Description	Parameters*
ARIMA	Autoregressive Integrated Moving Average	An automatic process to determine the differencing, autoregression, and moving averages that best fits historical time series data.	
DT ETS	Dampened Trend Exponential Smoothing	Weighted averages of historical time series data that are dampened through an exponentially decay function.	
Linear Trend	Linear Trend Model	Linear regression using historical time series data is used to calculate a trend.	
KKNN	K-Nearest Neighbor	The average of the “k” closet datapoints from the historical time series data.	K = 5, Distance Measure = “Minkowski Distance”, Weighting = “Optimal”
NNAR	Neural Network Autoregression	A feed-forward neural network where the trend is calculated from historical time series data.	Hidden Units = 10, Lagged Input = 1, Number of Networks = 20, Epochs = 100
XGBOOST	Extreme Gradient Boosting	An ensemble of regression tree algorithms that are gradient boosted.	Learning Rate = .3, Trees = 15, Maximum Tree Depth = 6
Drift*	Drift Method Benchmark	The average change from the last value of the historical time series data.	
Naïve*	Naïve Method Benchmark	The value of the most recent time series observation.	
SNaïve*	Seasonal Naïve Method Benchmark	The value of the most recent time series observation from the previous seasonal period.	
Mean*	Mean Benchmark	Uses the average value of historical time series data to forecast	

*Parameter defaults were used on all machine learning models. Those defaults are included in this column if appropriate.

Holdout: To evaluate forecast accuracy, we held out the latest two years of data from the time series (Williams & Kavanagh, 2016). Holding out data allows us to evaluate the accuracy of a forecast as we can compare the forecasted values against the holdout/actual data. Those forecasts that more closely match the holdout values will have smaller MAPE values. Our holdout data included the latest 24 months, 8 quarters, or 4 years of data. We doubled the recommendation on the annual data to increase the amount of holdout annual data.

Run Multiple Forecasts

We employed 10 different forecasting models: three classical approaches, three machine learning approaches, and four benchmarking methods. Three established methodologies were employed: the autoregressive integrated moving average (ARIMA), dampened trend exponential smoothing (DT ETS), and the linear trend model. Complementing these classic methods were three machine-learning algorithms: K-nearest neighbor (KNN), neural network autoregressive (NNAR), and extreme gradient boosting (XGBoost). Four benchmarking techniques were employed to provide baseline predictions: drift, naïve, seasonal naïve (SNaïve), and mean methods. A description of each model is included in Table 1.

Evaluate Against Benchmarks

MAPE scores were calculated and used in all evaluations. After running every revenue forecasting model, the MAPE scores were compared to those of the benchmark methods. The benchmark methods served an essential purpose, ensuring that a baseline level of performance was achieved.

Evaluate Overall Forecasting Accuracy

Outlier forecasts were also identified in addition to benchmark models. To identify outliers, MAPE scores were aggregated for every MP combination. Outliers were flagged when they were greater than three times the interquartile range. While this process is not practical for practitioners who only have access to their sales tax data, it is a consistency check in this study. Models that produce large numbers of outliers pose additional risks when used in practice.

Finally, we ranked each MP combination within each city to determine the relative accuracy. This provided a unique and often overlooked perspective in local government forecasting. Most studies measure model performance across time-series units, like Larson and Overton (2024). Therefore, to adhere to our city-centric approach to forecasting, we ranked at the city level rather than the time-series unit.

PREE: Results

The results of our analysis are presented using the PREE framework to illustrate how they inform city-level decisions.

Prepare

We present the visualized results using ridgeline plots that illustrate the relative MAPE ranking of each MP for each city by mapping the ranking density curve, revealing interesting patterns. To help readers interpret the graphs, note that “rank within city” reflects the relative performance of the forecast models for each city, with a rank closer to one indicating a more accurate forecast. For example, consider three hypothetical forecasting methods—Method 1, Method 2, and Method 3—applied to two cities, called City A and City B, with the following MAPE scores: for City A, Method 1 scores 0.04, Method 2 scores 1.7, and Method 3 scores 2.45; for City B, Method 1 scores 9.71, Method 2 scores 1.33, and Method 3 scores 4.25. Based on these scores, Method 1 would be ranked 1 for City A and 3 for City B; Method 2, ranked 2 for City A and 1

Figure 1. Preprocessing Rank Comparison Monthly Forecasts

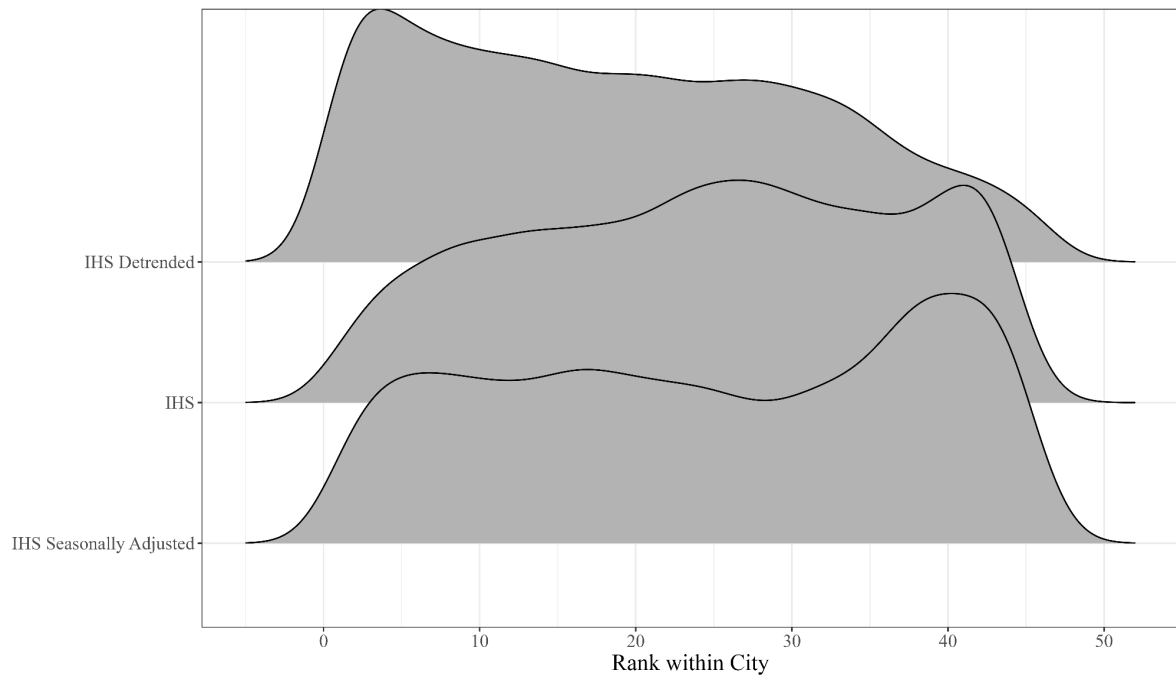


Figure 2. Preprocessing Rank Comparison Quarterly Forecasts

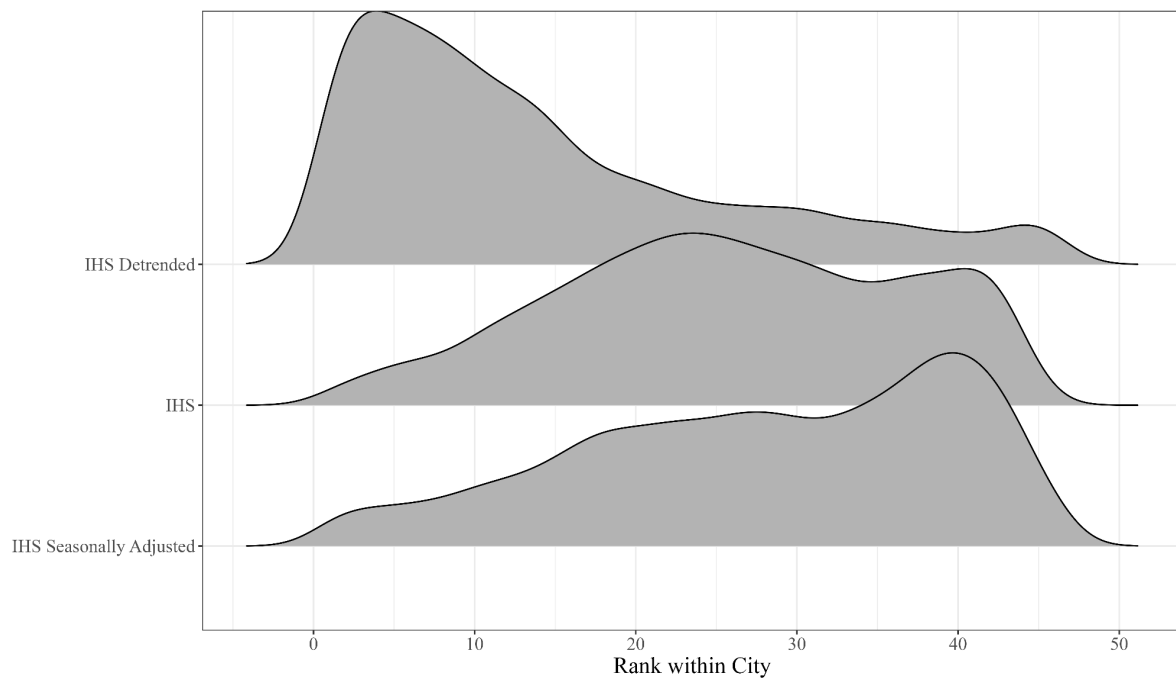


Figure 3. Model Rank Comparison Monthly Forecasts

Model Rank Comparison
Monthly Forecasts

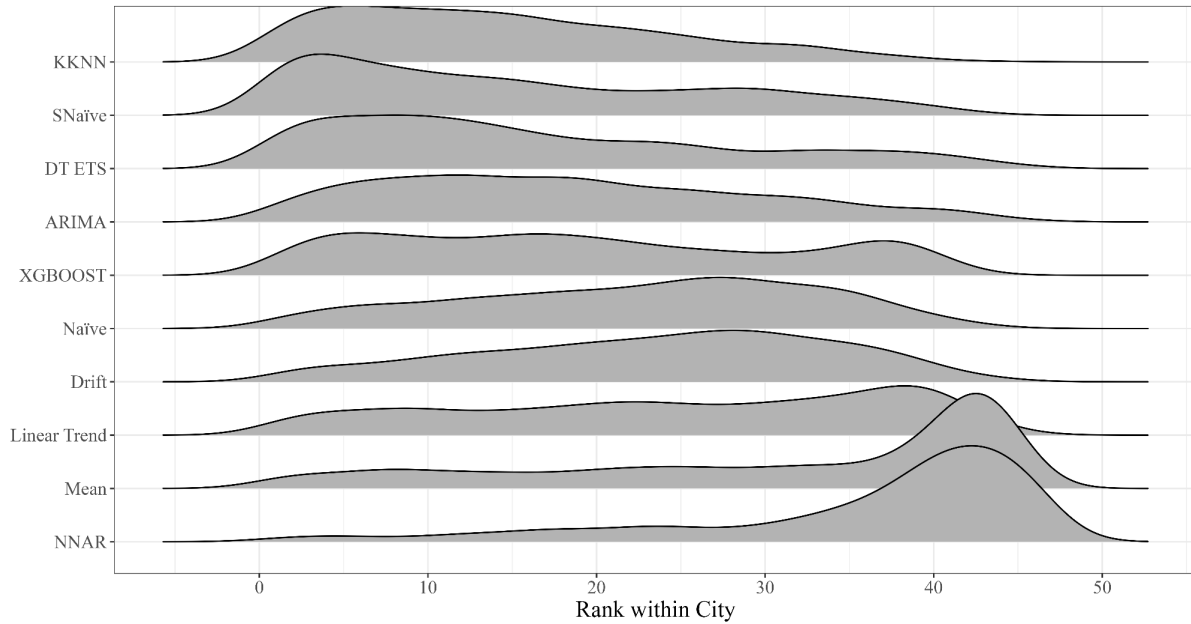


Figure 4. Model Rank Comparison Monthly Forecasts

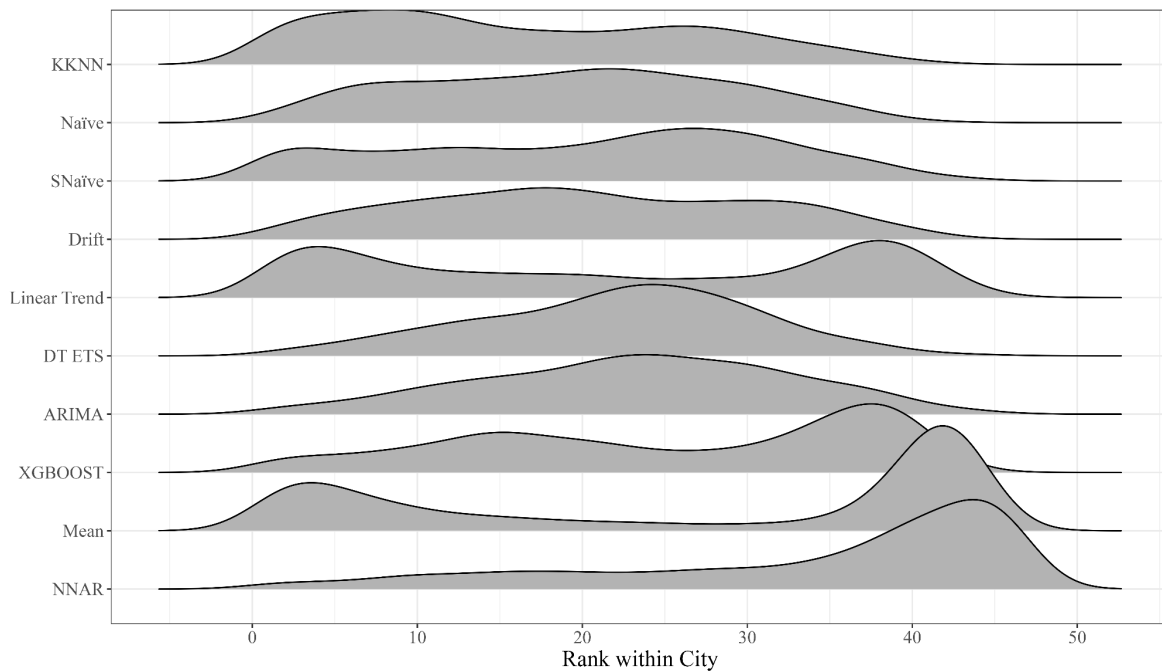
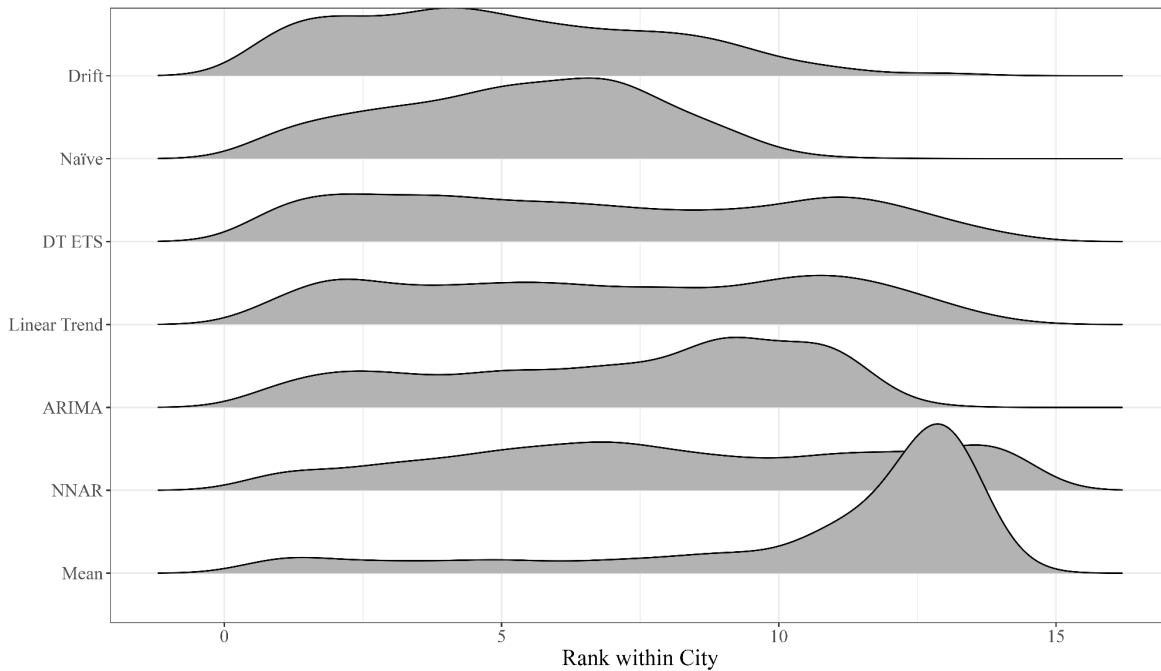


Figure 5. Model Rank Comparison Yearly Forecasts



for City B; and Method 3, ranked 3 for City A and 1 for City B. These rankings are then represented on the ridgeline plots, not the MAPE scores, enabling a clear comparison of the relative performance of a preprocessing or forecasting technique. To continue our example, we would create separate ridgeline plots for Methods 1, 2, and 3 to facilitate comparison of the ranking performance of a single method across cities, which is what we did for both preprocessing and forecasting methods in our analysis.

Higher density (i.e., larger hills) closer to the zero value indicates a higher number of cities where the forecasting technique performed the best compared to all others in a ranking system. The area under the grey curve represents all of the rankings that correspond with the x-axis values (titled “rank within city”). These plots provide a quick way to compare all forecast rankings across models and preprocessing procedures. Therefore, according to that forecasting technique, the higher the peak, the more cities experienced the same ranking.

For Figures 1 and 2, the rankings of each forecast are aggregated to the preprocessing step, which means that within each preprocessing step, the rank of all the different models is combined with that preprocessing step. Preprocessing step comparisons for monthly (Figure 1) and quarterly data (Figure 2) indicated that each preprocessing variant could produce the most accurate forecasts for a city, with IHS detrended consistently emerging as the preprocessing approach used in high-ranking forecasts for both monthly and quarterly data. While less consistent than models using IHS detrended preprocessing, IHS and IHS Seasonally Adjusted preprocessing steps also produced top-ranking forecasts for some cities in our sample, reinforcing the need for evaluation to determine each city’s best-performing version. This finding suggests that PREE is necessary for preprocessing since all preprocessing steps perform well in some cities.

Table 2. Benchmark and Outlier Comparison

	At Least One Benchmark Model Outperforms			
	All Forecasts	Some, But Not All Forecasts	No Other Forecasts	Total
Month				
Not an Outlier	1,573 (17.4%)	5,413 (59.9%)	1,324 (14.6%)	8,310 (91.9%)
Outlier	63 (0.7%)	397 (4.4%)	272 (3.0%)	732 (8.1%)
Quarter				
Not an Outlier	1,637 (15.2%)	6,171 (57.5%)	2,080 (19.4%)	9,888 (92.1%)
Outlier	85 (0.8%)	434 (4.0%)	329 (3.1%)	848 (7.9%)
Year				
Not an Outlier	794 (19.8%)	2,373 (59.0%)	570 (14.2%)	3,737 (93.0%)
Outlier	9 (0.2%)	92 (2.3%)	182 (4.5%)	283 (7.0%)

Note: Includes mean, drift, and naïve benchmark models for all periods and seasonal naïve models for monthly and quarterly data.

Run Multiple Forecasts

Similar to the comparison of preprocessing steps, ridgeline plots present model forecasting ranks for monthly (Figure 3), quarterly (Figure 4), and yearly (Figure 5) data, where the rankings of forecasts using different preprocessing steps are aggregated for each model. Figures 3 to 5 are presented with the most effective approach at the top, followed by the remaining approaches in order of effectiveness. These plots enable comparisons of relative model rankings. While some models, such as KNN, consistently ranked high relative to other models, and some, like NNAR and the mean benchmark model, consistently ranked low, *all models produced the most accurate forecast for at least one city in the sample*. Conversely, all models produced forecasts ranking outside the top 10 most accurate models for at least one city. This variability highlights the risk of relying on a single gold standard model for local government forecasting. PREE advocates for evaluating multiple forecasting methods, and in the absence of a consistently top-performing approach, cities must implement a framework to assess and compare these methods. It emphasizes running and comparing multiple models to select the best-performing approach.

Evaluate Against Benchmarks

The city-centric analysis in Table 2 revealed that only a tiny percentage of model or preprocessing forecasts outperformed all benchmark models while avoiding outlier status across all cities in the sample. Outliers were identified as MAPE values three times the interquartile range of each model/preprocessing combination. Specifically, 17.3%, 15.2%, and 19.8% of forecasts met these monthly, quarterly, and yearly data criteria, respectively. Notably, 7–8% of the forecasts produced outlying forecasts. Fewer than 1% of models across all periods were outliers, yet they still outperformed all four benchmark models. These findings underscore the utility of benchmark models as a practical tool for individual forecasters to evaluate the efficacy of various forecasting methods without requiring comparisons to other cities.

Evaluate Overall Forecasting Accuracy

Table 3. Within City Rank Comparison, Monthly Forecasts

	Rank					Outlier MAPE				
	1	2	3	Avg.	Outliers	Avg.	Min	Med.	Max	Avg.
SNaïve-IHS Detrended	163	116	76	8.25	62	3.96	2.73	3.45	8.29	1.13
KKNN-IHS Detrended	63	39	47	11.35	63	3.41	2.35	3.14	6.71	1.19
XGBOOST-IHS Detrended	29	28	40	14.07	66	4.23	2.47	3.45	22.36	1.30
DT ETS-IHS	44	24	35	16.40	63	6.20	3.46	4.84	33.73	1.60
KKNN-IHS Seasonally	16	15	23	17.79	62	5.41	3.40	4.44	31.00	1.57
Adjusted										
ARIMA-IHS	16	17	21	18.48	75	6.54	3.32	4.12	114.36	1.71
SNaïve-IHS Seasonally	5	4	37	18.84	64	5.64	3.56	4.68	30.43	1.61
Adjusted										
Naïve-IHS Seasonally	26	20	32	18.92	63	6.21	3.80	4.93	22.71	1.68
Adjusted										
SNaïve-IHS	2	7	13	19.09	64	5.65	3.55	4.69	30.37	1.61
Linear Trend-IHS Detrended	13	17	29	19.34	75	3.26	2.34	2.91	6.73	1.38
Mean-IHS Detrended	27	16	27	19.62	74	3.27	2.35	2.92	6.78	1.38
Drift-IHS Seasonally	11	30	21	22.07	66	6.38	4.00	5.09	23.96	1.81
Adjusted										
Naïve-IHS Detrended	4	4	8	22.47	80	4.86	2.75	3.86	22.47	1.64
Drift-IHS Detrended	3	2	6	23.23	87	4.77	2.71	3.82	23.21	1.66
Linear Trend-IHS Seasonally	25	23	15	24.52	46	8.03	5.02	6.74	34.98	2.10
Adjusted										
XGBOOST-IHS Seasonally	5	7	19	24.72	69	5.36	3.41	4.55	22.55	1.79
Adjusted										
Drift-IHS	1	2	2	25.35	89	5.47	3.17	4.38	24.26	1.83
Naïve-IHS	0	2	1	26.22	82	5.56	3.18	4.41	23.29	1.84
Linear Trend-IHS	2	0	5	29.93	72	6.83	4.37	5.27	34.93	2.26
NNAR-IHS Detrended	7	3	10	35.03	66	18.39	8.93	12.04	91.79	4.02
Mean-IHS Seasonally	15	2	7	35.32	27	15.26	11.80	14.20	38.74	4.12
Adjusted										
NNAR-IHS Seasonally	3	0	1	36.35	75	29.80	10.53	15.54	285.65	5.43
Adjusted										
Mean-IHS	2	2	1	37.09	28	15.12	11.53	14.17	38.66	4.19

Further analysis using rank comparison (Tables 3 through 5) provided a more nuanced look at relative MP performance within each city. The far left column articulates the MP combination where the first word is the model used, followed by the preprocessing steps. The second, third, and fourth columns represent the counts of cities where the MP combination achieved this relative ranking. The fifth column is the average rank of the MP combination across cities. The “Outliers” column represents the number of cities where this MP combination yielded an outlying forecast. The remaining columns show the MAPE summary statistics for each MP combination.

Table 4. Within City Rank Comparison, Quarterly Forecasts

	Rank				Outliers	Outlier MAPE				
	1	2	3	Avg.		Avg.	Min	Med.	Max	Avg.
Mean-IHS Detrended	109	117	116	6.55	81	2.06	1.24	1.75	7.16	0.54
Linear Trend-IHS Detrended	92	101	127	6.72	83	2.06	1.24	1.75	7.17	0.55
KKNN-IHS Detrended	85	64	40	9.56	76	2.33	1.49	1.97	7.33	0.62
SNaïve-IHS Detrended	94	112	45	11.36	65	3.03	1.88	2.60	12.13	0.71
Naïve-IHS Detrended	8	16	21	14.62	82	3.35	1.89	2.63	12.73	0.83
Drift-IHS Detrended	15	9	24	15.68	83	3.42	1.92	2.65	13.14	0.85
XGBOOST-IHS Detrended	24	19	25	19.40	136	2.19	1.20	1.87	12.84	0.83
Naïve-IHS Seasonally Adjusted	16	8	19	21.22	61	5.31	3.07	4.18	16.94	1.19
Naïve-IHS	4	3	8	21.96	66	5.19	2.95	4.18	16.89	1.20
DT ETS-IHS	2	6	9	22.08	67	4.95	2.85	4.10	20.43	1.18
Drift-IHS Seasonally Adjusted	12	14	13	22.63	64	5.34	3.11	4.35	16.54	1.22
Drift-IHS	8	6	6	22.87	68	5.24	3.03	4.29	16.86	1.23
KKNN-IHS Seasonally Adjusted	10	8	9	23.04	70	5.01	3.12	4.32	25.15	1.27
ARIMA-IHS	7	6	6	23.23	74	5.42	2.98	4.20	27.43	1.28
SNaïve-IHS Seasonally Adjusted	3	6	9	24.85	75	5.13	3.31	4.19	25.07	1.33
SNaïve-IHS	2	4	9	24.87	74	5.16	3.36	4.25	25.05	1.33
Linear Trend-IHS Seasonally Adjusted	14	6	4	29.06	45	9.03	5.06	7.15	31.07	1.88
Linear Trend-IHS	4	2	2	29.49	46	9.00	5.02	7.17	31.10	1.89
XGBOOST-IHS Seasonally Adjusted	6	5	7	31.10	72	4.96	3.17	4.27	16.25	1.67
NNAR-IHS Detrended	18	5	10	31.61	83	33.49	11.88	19.18	475.28	5.14
NNAR-IHS Seasonally Adjusted	2	0	3	35.96	96	58.69	14.83	24.52	1182.64	8.65
Mean-IHS Seasonally Adjusted	11	1	3	36.97	32	15.24	12.01	14.12	31.87	3.90
Mean-IHS	3	5	1	37.00	33	15.13	11.56	14.10	31.84	3.91

Monthly data are in Table 3, quarterly data in Table 4, and yearly data in Table 5. This analysis revealed some top-performing forecasts that were not apparent when examining the model ridgelines graphs. For instance, while mean models generally ranked low for the quarterly forecasts, the mean-IHS detrended forecast produced the most top-ranked forecasts for the quarterly data. Importantly, all MP combinations produced a non-trivial number of city outlier forecasts, ranging from 28 to 136 across monthly, quarterly, and yearly forecasts. Even the best MP combinations generated outlying forecasts for some cities.

Table 5. Within City Rank Comparison, Yearly Forecasts

	Rank				Outliers	Outlier MAPE				
	1	2	3	Avg.		Avg.	Min	Med.	Max	Avg.
Drift	112	122	104	4.99	63	6.40	3.80	4.66	20.72	1.40
Naïve	64	83	102	5.31	53	6.54	3.85	5.40	17.74	1.45
DT ETS	83	97	87	6.69	63	10.85	5.97	8.80	41.36	2.08
Linear Trend	56	105	73	6.98	54	8.43	5.30	6.79	28.45	1.88
ARIMA	51	75	70	7.09	63	6.56	4.50	5.75	16.75	1.76
NNAR	46	36	56	8.18	103	30.21	8.95	14.08	638.20	4.89
Mean	41	25	24	10.51	36	12.83	10.43	12.23	19.97	3.63

Further analysis examining the prevalence of outlier forecasts among cities is presented using density plots for the monthly (Figure 6), quarterly (Figure 7), and yearly (Figure 8) data. The outlier comparison revealed that approximately one-third of cities were outliers in at least one forecasting model: 33% (271 unique cities) for monthly forecasts, 31% (303 unique cities) for quarterly forecasts, and 20.5% (206 unique cities) for annual forecasts. Only a few cities were outliers in most of the forecasting models (MP) combinations used in this study. Most cities that produced an outlying forecast only generated one outlier. These distributions indicate that a select few cities do not drive outlier counts but rather are spread across the entire sample.

These results collectively emphasize that no single approach to model selection in sales tax revenue forecasting is universally applicable. The variability in model performance across cities underscores the importance of running various models, comparing their forecasting accuracy, and carefully selecting the best-performing models for each specific municipality.

Discussion

We used a city-centric approach to understand sales tax revenue forecasting, which challenged the notion of a universal gold standard forecasting method for local governments. The results underscore the complexity and variability inherent in municipal-level sales tax revenue forecasting, as well as how city-centric approaches reveal the risks practitioners face when engaging with revenue forecasting scholarship. The findings broadly suggested that practitioners should adopt PREE rather than relying solely on the latest or most complex forecasting method. Such processes encourage practitioners to use multiple models, integrate benchmarking, and make evaluation and comparison a routine aspect of forecasting. No single model that we have evaluated justifies a one-size-fits-all approach. This study's findings challenge conventional wisdom about revenue forecasting and highlight several key insights that underscore the importance of a more nuanced, city-centric approach to forecasting methodology.

First, only a small percentage (15–20%) of forecasts in our analysis outperformed all benchmark models while avoiding outlier status across monthly, quarterly, and yearly data. This finding highlights the difficulty of consistently producing accurate forecasts with a single modeling approach. Benchmark models offer several benefits, including the ability to identify

Figure 6. Monthly Forecasts

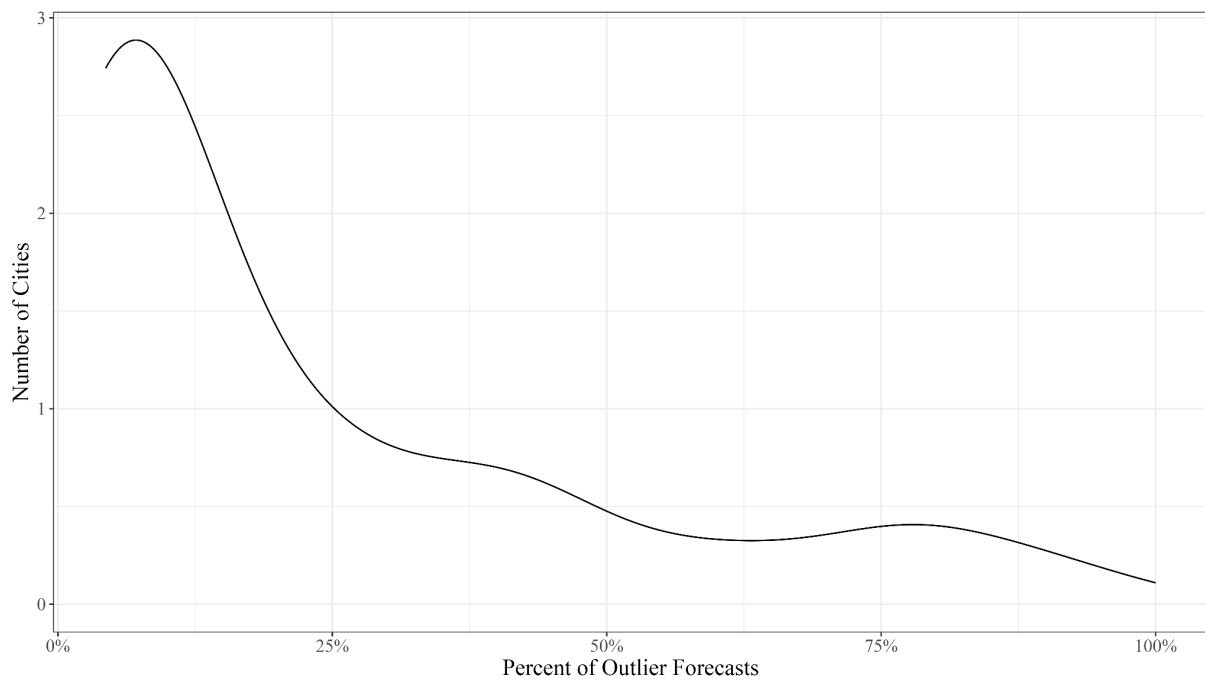


Figure 7. Quarterly Forecasts

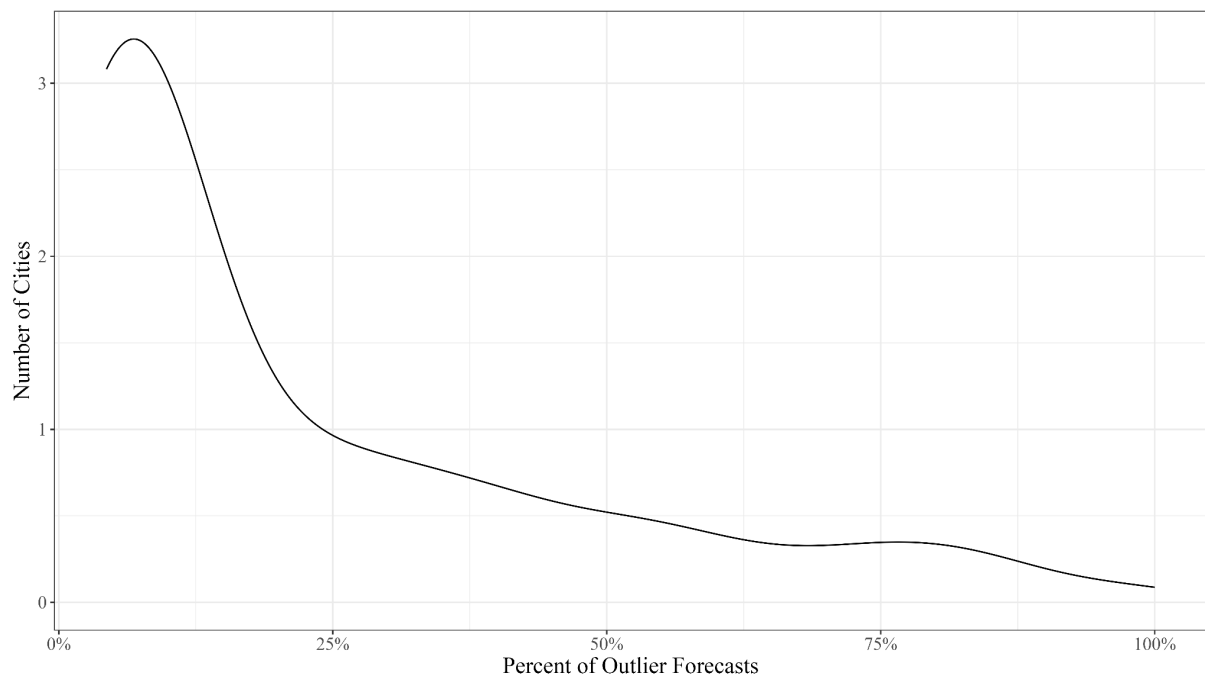
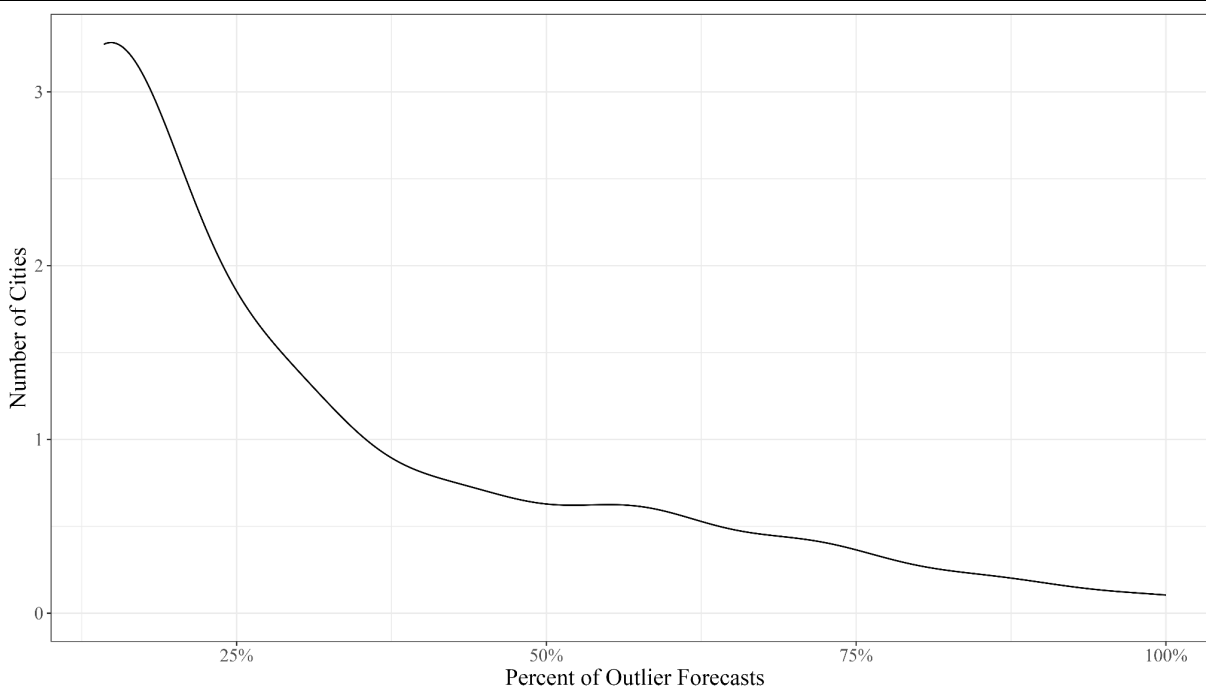


Figure 8. Yearly Forecasts



poorly performing forecasts early in the evaluation process. They can be quickly adopted as a tool for individual forecasters to evaluate the efficacy of various forecasting methods.

Second, the variability in relative model and preprocessing step performance, as illustrated by the ridgeline plots, is particularly striking since all models were the most accurate forecast for at least *one city* while simultaneously producing forecasts outside the top 10 most accurate for other cities. While IHS preprocessing emerged as a consistently strong performer, the variability in performance across other steps, such as detrending, underscores the need for careful evaluation to determine the best-performing version for each specific municipality. Cities produce highly localized revenue patterns that can make certain generally highly accurate models and preprocessing steps inaccurate. The safest way to minimize risk is to have a robust forecasting evaluation process in place.

Third, the outlier analysis offers further insights into the risks associated with relying on a single forecasting method. The fact that all MP combinations generated outlying forecasts for some cities, and that approximately one-third of cities generated at least one outlying forecast, highlights the potential for significant forecasting errors if a diverse range of methods is not considered.

For researchers, these results call for a shift in focus from identifying universally optimal forecasting methods to developing frameworks that can guide practitioners in selecting and evaluating methods tailored to their specific contexts. For scholars interested in studying forecasting accuracy, assessing relative accuracy using city-level rankings can help frame the results meaningfully for practitioners. Furthermore, tax forecasting accuracy scholars should strongly consider the routine inclusion of benchmark models to ensure that models generally exceed minimum accuracy thresholds, and outlier analysis to help account for the risk of generating relatively inaccurate forecasts for cities using such methods.

Conclusion

This study challenges the notion of a gold standard forecasting method. It highlights the need for a more nuanced and localized perspective on sales tax revenue forecasting in local governments. In many ways, this study's findings suggest that the focus of a vast majority of revenue forecasting academic research has not been particularly useful to practitioners, as it has been largely driven by the scholarly pursuit of finding the forecasting gold standard. By examining forecasting accuracy at the individual city level, we have demonstrated the high variability in model performance across municipalities and the associated risks of relying on a single forecasting approach.

The PREE framework proposed in this study offers a practical guide for local government practitioners to navigate the complexities of sales tax revenue forecasting. This approach emphasizes the importance of employing multiple models, using benchmarks, and implementing regular evaluation standards to help practitioners achieve more accurate and reliable forecasts.

Future research should focus on developing more sophisticated frameworks for matching forecasting methods to specific municipal characteristics, exploring the factors that drive variability in model performance across cities, and investigating how changing economic conditions impact the relative performance of different forecasting approaches. These suggestions highlight the need to shift away from the existing literature's search for a gold standard model and towards a robust forecasting process. Future research should build upon the work of Kriz (2019), exploring the potential of ensemble forecasting approaches to achieve increased efficiency gains, with a focus on the utility of ensemble forecasting at a city-centric level.

The results of this study suggest that scholars can help practitioners by focusing on the unique features that would make particular techniques appropriate for inclusion within the PREE framework for practitioners. Additionally, future research should include studies that identify other data sources to help improve the accuracy of time-series forecasts. Evaluation of ensemble approaches is needed to further the science and practice of local government forecasting. Regardless of the specific focus of future research, we encourage a focus on the practical application for practitioners.

Disclosure Statement

The authors declare no conflicts of interest related to this article's research, authorship, or publication.

References

- Bretschneider, S. I., & Gorr, W. L. (1992). Economic, organizational, and political influences on biases in forecasting state sales tax receipts. *International Journal of Forecasting*, 7(4), 457-466. [https://doi.org/10.1016/0169-2070\(92\)90029-9](https://doi.org/10.1016/0169-2070(92)90029-9)

- Bretschneider, S. I., Gorr, W. L., Grizzle, G., & Klay, E. (1989). Political and organizational influences on the accuracy of forecasting state government revenues. *International Journal of Forecasting*, 5(3), 307-319. [https://doi.org/10.1016/0169-2070\(89\)90035-6](https://doi.org/10.1016/0169-2070(89)90035-6)
- Brogan, M. (2012). The politics of budgeting: Evaluating the effects of the political election cycle on state-level budget forecast errors. *Public Administration Quarterly*, 36(1), 84-115.
- Buettner, T., & Kauder, B. (2010). Revenue forecasting practices: Differences across countries and consequences for forecasting performance. *Fiscal Studies*, 31(3), 313-340. <https://doi.org/10.1111/j.1475-5890.2010.00117.x>
- Buxton, E., Kriz, K., Cremeens, M., & Jay, K. (2019). An auto regressive deep learning model for sales tax forecasting from multiple short time series. In *2019 18th IEEE international conference on machine learning and applications (ICMLA)* (pp. 1359–1364). IEEE. <https://doi.org/10.1109/ICMLA.2019.00221>
- Caiden, N., & Wildavsky, A. B. (1980). *Planning and budgeting in poor countries*. Transaction Publishers
- Carmody, K., & Wiipongwii, T. (2018). *Virginia tax revenue forecasting* [Unpublished Manuscript]. Department of Public Policy, College of William and Mary.
- Cassar, G., & Gibson, B. (2008). Budgets, internal reports, and manager forecast accuracy. *Contemporary Accounting Research*, 25(3), 707-738. <https://doi.org/10.1506/car.25.3.3>
- Chapman, J. I. (1982). Fiscal stress and budgetary activity. *Public Budgeting & Finance*, 2(2), 83-87. <https://doi.org/10.1111/1540-5850.00561>
- Chung, I. H., Williams, D. W., & Do, M. R. (2022). For better or worse? Revenue forecasting with machine learning approaches. *Public Performance & Management Review*, 45(5), 1133-1154. <https://doi.org/10.1080/15309576.2022.2073551>
- Cover, T., & Hart, P. (1967). Nearest neighbor pattern classification. *IEEE Transactions on Information Theory*, 13(1), 21-27. <https://doi.org/10.1109/TIT.1967.1053964>
- Deschamps, E. (2004). The impact of institutional change on forecast accuracy: A case study of budget forecasting in Washington State. *International Journal of Forecasting*, 20(4), 647-657. <https://doi.org/10.1016/j.ijforecast.2003.11.009>
- Forrester, J. P. (1991). Budgetary constraints and municipal revenue forecasting. *Policy Sciences*, 24(4), 333-356. <https://doi.org/10.1007/BF00135880>
- Frank, H. A., & Zhao, Y. (2009). Determinants of local government revenue forecasting practice: Empirical evidence from Florida. *Journal of Public Budgeting, Accounting & Financial Management*, 21(1), 17-35. <https://doi.org/10.1108/JPBAFM-21-01-2009-B002>
- Fullerton, T. M., Jr. (1989). A composite approach to forecasting state government revenues: Case study of the Idaho sales tax. *International Journal of Forecasting*, 5(3), 373-380. [https://doi.org/10.1016/0169-2070\(89\)90040-X](https://doi.org/10.1016/0169-2070(89)90040-X)
- Hansen, J. V., & Nelson, R. D. (1997). Neural networks and traditional time series methods: A synergistic combination in state economic forecasts. *IEEE Transactions on Neural Networks*, 8(4), 863-873. <https://doi.org/10.1109/72.595884>
- Hansen, J. V., & Nelson, R. D. (2002). Data mining of time series using stacked generalizers. *Neurocomputing*, 43(1-4), 173-184. [https://doi.org/10.1016/S0925-2312\(00\)00364-7](https://doi.org/10.1016/S0925-2312(00)00364-7)
- Hyndman, R. J., & Athanasopoulos, G. (2018). *Forecasting: Principles and practice*. OTexts.
- Kong, D. (2007). Local government revenue forecasting: The California county experience. *Journal of Public Budgeting, Accounting & Financial Management*, 19(2), 178-199. <https://doi.org/10.1108/JPBAFM-19-02-2007-B003>

- Krause, P. (2012). Executive politics and the governance of public finance. In M. Lodge & K. Wegrich (Eds.), *Executive politics in times of crisis* (pp. 136-156). Palgrave Macmillan. https://doi.org/10.1057/9781137010261_8
- Kriz, K. A. (2019). Ensemble Forecasting. In D. W. Williams & T. D. Calabrese (Eds.), *The Palgrave handbook of government budget forecasting* (pp. 413–426). Palgrave Macmillan. https://doi.org/10.1007/978-3-030-18195-6_21
- Larson, S., & Overton, M. (2024). Modeling approach matters, but not as much as preprocessing: Comparison of machine learning and traditional revenue forecasting techniques. *Public Finance Journal*, 1(1), 29-48. <https://doi.org/10.59469/pfj.2024.8>
- Lorenz, T., & Homburg, C. (2018). Determinants of analysts' revenue forecast accuracy. *Review of Quantitative Finance and Accounting*, 51, 389-431. <https://doi.org/10.1007/s11156-017-0675-4>
- MacManus, S. A., & Grothe, B. P. (1989). Fiscal stress as a stimulant to better revenue forecasting and productivity. *Public Productivity Review*, 12(4), 387-400. <https://doi.org/10.2307/3380152>
- Makridakis, S., Hyndman, R. J., & Petropoulos, F. (2020). Forecasting in social settings: The state of the art. *International Journal of Forecasting*, 36(1), 15-28. <https://doi.org/10.1016/j.ijforecast.2019.05.011>
- Mikesell, J. L. (2018). Often wrong, never uncertain: Lessons from 40 years of state revenue forecasting. *Public Administration Review*, 78(5), 795-802. <https://doi.org/10.1111/puar.12954>
- Mikesell, J. L., & Ross, J. M. (2014). State revenue forecasts and political acceptance: The value of consensus forecasting in the budget process. *Public Administration Review*, 74(2), 188-203. <https://doi.org/10.1111/puar.12166>
- Mocan, H. N., & Azad, S. (1995). Accuracy and rationality of state general fund revenue forecasts: Evidence from panel data. *International Journal of Forecasting*, 11(3), 417-427. [https://doi.org/10.1016/0169-2070\(95\)00592-9](https://doi.org/10.1016/0169-2070(95)00592-9)
- Muh, K. L., & Jang, S. B. (2019). A time-series data prediction using tensorflow neural network libraries. *KIPS Transactions on Computer and Communication Systems*, 8(4), 79-86. <https://doi.org/10.3745/KTCCS.2019.8.4.79>
- Reddick, C. G. (2004). Assessing local government revenue forecasting techniques. *International Journal of Public Administration*, 27(8-9), 597-613. <https://doi.org/10.1081/PAD-120030257>
- Reitano, V. (2018). An open systems model of local government forecasting. *American Review of Public Administration*, 48(5), 476-489. <https://doi.org/10.1177/0275074017692876>
- Rodgers, R., & Joyce, P. (1996). The effect of underforecasting on the accuracy of revenue forecasts by state governments. *Public Administration Review*, 56(1), 48-56. <https://doi.org/10.2307/3110053>
- Rubin, I. S. (1987). Estimated and actual urban revenues: Exploring the gap. *Public Budgeting & Finance*, 7(4), 83-94. <https://doi.org/10.1111/1540-5850.00766>
- Rubin, M. M., Mantell, N., & Pagano, M. A. (1999). Approaches to revenue forecasting by state and local governments. *Proceedings. Annual Conference on Taxation and Minutes of the Annual Meeting of the National Tax Association*, 92, 205-221.
- Rubin, M. M., Peters, J. L., & Mantell, N. (2019). Revenue forecasting and estimation. In W. B. Hildreth (Ed.), *Handbook on taxation* (pp. 769-800). Routledge.

- Texas State Comptroller. (2022). *Transparency: Local government*.
<https://comptroller.texas.gov/transparency/local/cities.php>
- Tversky, A., & Kahneman, D. (1992). Advances in prospect theory: Cumulative representation of uncertainty. *Journal of Risk and Uncertainty*, 5, 297-323.
<https://doi.org/10.1007/BF00122574>
- Voorhees, W. R. (2006). Neural networks and revenue forecasting: A smarter forecast? *International Journal of Public Policy*, 1(4), 379-388.
<https://doi.org/10.1504/IJPP.2006.010843>
- Williams, D. W. (2012). The politics of forecast bias: Forecaster effect and other effects in New York City revenue forecasting. *Public Budgeting & Finance*, 32(4), 1-18.
<https://doi.org/10.1111/j.1540-5850.2012.01021.x>
- Williams, D. W., & Calabrese, T. D. (2016). The status of budget forecasting. *Journal of Public and Nonprofit Affairs*, 2(2), 127-160. <https://doi.org/10.20899/jpna.2.2.127-160>
- Williams, D. W., & Kavanagh, S. C. (2016). Local government revenue forecasting methods: Competition and comparison. *Journal of Public Budgeting, Accounting & Financial Management*, 28(4), 488-526. <https://doi.org/10.1108/JPBAFM-28-04-2016-B004>

Author Biographies

Sarah E. Larson is an associate professor in the Department of Political Science at Miami University. She also serves as the social media editor for both *Public Finance Journal* and *Public Administration*, as well as the treasurer of the Association for Budgeting and Financial Management. She received her B.A. in political science and government from Miami University, her MPA, M.S. in statistics, and Ph.D. in public policy analysis from Indiana University. Her research focuses on issues in taxation policy; specifically, she examines sales, use, income, and property taxation and methodological advancements in public administration and policy.

Michael R. Overton is an associate professor in the Department of Politics and Philosophy and the associate director of the Institute for Interdisciplinary Data Sciences at the University of Idaho. He received his B.A. in music, MPA, and Ph.D. in public administration from the University of North Texas. His research focuses on the integration of data science, public finance, and public administration through data science literacy and administrative informatics.