
ESTIMATING NATIONWIDE HIGH-DOSAGE TUTORING EXPENDITURES: A PREDICTIVE MODEL APPROACH

A PREPRINT

 **Jason Godfrey**
Strategic Data Fellow
Harvard University
Boston, MA 02138

 **Trisha Banerjee**
Data Engineer
Amazon Web Services

August 9, 2024

ABSTRACT

This study applies an optimized XGBoost regression model to estimate district-level expenditures on high-dosage tutoring from incomplete administrative data. The COVID-19 pandemic caused unprecedented learning loss, with K-12 students losing up to half a grade level in certain subjects. To address this, the federal government allocated \$190 billion in relief. We know from previous research that small-group tutoring, summer and after school programs, and increased support staff were all common expenditures for districts. We don't know how much was spent in each category. Using a custom scraped dataset of over 7,000 ESSER (Elementary and Secondary School Emergency Relief) plans, we model tutoring allocations as a function of district characteristics such as enrollment, total ESSER funding, urbanicity, and school count. Extending the trained model to districts that mention tutoring but omit cost information yields an estimated aggregate allocation of approximately \$2.2 billion. The model achieved an out-of-sample $R^2=0.358$, demonstrating moderate predictive accuracy given substantial reporting heterogeneity. Methodologically, this work illustrates how gradient-boosted decision trees can reconstruct large-scale fiscal patterns where structured data are sparse or missing. The framework generalizes to other domains where policy evaluation depends on recovering latent financial or behavioral variables from semi-structured text and sparse administrative sources.

Keywords Education Policy · ESSER · Predictive Modeling · XGBoost · Educational Finance

1 Introduction

The COVID-19 pandemic has profoundly impacted the educational landscape, exacerbating existing inequalities and resulting in significant learning loss among K-12 students. This loss, quantified as up to half a grade level in some subjects, has prompted a vigorous response from educators and policymakers aiming to mitigate the long-term consequences of this disruption [Kuhfeld et al., 2020, Dorn et al., 2020]. In response, the federal government allocated approximately \$190 billion in relief funding through the Elementary and Secondary School Emergency Relief (ESSER) funds, a substantial investment aimed at addressing these educational deficits.

Among the various interventions proposed, high-dosage tutoring has emerged as one of the most promising strategies for mitigating learning loss [Nickow et al., 2020]. High-dosage tutoring, characterized by intensive, small-group or one-on-one instruction, has been shown to consistently positively boost student achievement, particularly in reading and mathematics [Kraft et al., 2022]. However, while the theoretical benefits of high-dosage tutoring are well-documented, less is known about the actual allocation of ESSER funds towards this intervention across different school districts.

This study seeks to bridge this knowledge gap by estimating district-level expenditures on high-dosage tutoring. Utilizing a custom dataset scraped from over 7,000 ESSER plans, we implement an optimized XGBoost regression model [Chen and Guestrin, 2016] to infer spending levels among districts that reference tutoring but do not report explicit dollar amounts. This approach treats financial inference as a predictive modeling problem—leveraging a broad

set of district-level features such as urbanicity, total ESSER allocation, and school count—to impute unobserved values. Beyond its substantive findings, the paper contributes methodologically by demonstrating how ensemble learning can reconstruct incomplete administrative data at national scale, offering a transparent, reproducible framework for policy analytics.

2 Literature Review

The literature on educational interventions post-pandemic underscores the critical need for effective resource allocation [Dewey et al., 2024]. High-dosage tutoring, characterized by frequent, intensive sessions, has been widely recognized as a potent strategy for addressing learning loss, particularly for students from disadvantaged backgrounds [Nickow et al., 2020]. Studies have consistently demonstrated the effectiveness of such tutoring in improving academic outcomes [Kraft et al., 2022, Robinson et al., 2021].

Previous work emphasizes the profound impacts of the pandemic on student achievement and the necessity of targeted interventions to bridge the resultant gaps [Carbonari et al., 2024]. Research by Gwynne and Allensworth [2023], Balfanz and Byrnes [2023] corroborates these findings, highlighting the efficacy of high-dosage tutoring in enhancing student performance in core subjects like reading and mathematics.

Recent meta-analyses reveal that small-group tutoring can significantly reduce achievement gaps, with the most substantial benefits observed in low-income student populations [Dietrichson et al., 2017]. These analyses show that interventions such as structured tutoring and feedback mechanisms are essential for improving educational outcomes among disadvantaged students. Additionally, studies indicate that the long-term economic benefits of improved educational outcomes far outweigh the initial costs of such programs, making high-dosage tutoring a cost-effective solution [Guryan et al., 2023].

Ensemble and supervised machine learning methods have increasingly been applied in educational, administrative, and policy-relevant settings to handle large, heterogeneous datasets and infer or predict outcomes where traditional econometric approaches may struggle. For example, a systematic review of predictive models in education finds that machine learning algorithms—including gradient boosting and tree-based methods—consistently outperform classical statistical models in forecasting student outcomes, managing non-linear relationships, and handling high-dimensional inputs [Almalawi et al., 2024]. In one study, the tree-boosting algorithm XGBoost achieved superior performance in predicting learner performance on seven datasets compared with Item Response Theory models and standard logistic regression frameworks [Hakkal and Ait Lahcen, 2024]. In policy- or finance-oriented contexts, gradient boosting methods (including XGBoost) have been shown to outperform alternatives in settings such as banking failure and financial distress modeling [Carmona et al., 2019, Lokanan and Ramzan, 2024]. These findings signal the utility of ensemble learning for inferring latent or unreported values in administrative data, such as resource allocations across large populations—a methodology closely aligned with the present study’s objective of imputing district-level tutoring expenditures. Collectively, this line of research supports treating fiscal inference as a predictive modeling task using large-scale administrative inputs, offering novel methodological possibilities for education policy analytics.

3 Data Sources

This study uses two custom data sources which both derive from the ESSER spending documents that each district was required to create in order to receive funding. The first dataset contains hand-scraped information from each of the ESSER documents regarding explicitly mentioned budget items such as tutoring, summer learning, and credit recovery. This dataset contains 7024 (roughly 36% of total districts) unique districts with a combined enrollment of 41.1 million students, or roughly 83.45% of estimated k12 enrollment in the USA. The second dataset contains 4685 district plans that have been passed through optical character recognition and their text can be searched. The intersection of these datasets is 4387 districts that can be analyzed for both the textual presence of "tutor*" and on specific expenditures that were manually extracted. This breakdown can be viewed in figure Figure 1.

Within the intersected dataset, 1232 districts defined how much ESSER money they allocated to tutoring. The mean tutoring allocation was \$263,213.68 with a standard deviation of \$296,770.74. This distribution can be seen in Figure 2. Note that this histogram is drawn after removing outliers, such as Houston ISD that allocated 113.33 million to "High Dosage Tutorials."

This dataset, while incomplete, provides the best available insight into district-level expenditures of ESSER funds on all items not recorded by the federal government, such as high-dosage tutoring. All these plans are publicly available on their respective district websites and they contain no sensitive information.

4 Methods

In this study, we employed the XGBoost (Extreme Gradient Boosting) algorithm to predict tutoring expenses based on various features. XGBoost is an efficient and scalable implementation of gradient boosting, a powerful machine learning technique for regression and classification problems.

4.1 Gradient Boosting Framework

Gradient boosting builds an ensemble of weak learners, typically decision trees, by sequentially fitting new models to the residual errors made by previous models. The objective is to minimize the loss function $L(y_i, \hat{y}_i)$, where y_i is the true value and $\hat{y}_i$ is the predicted value. Given a training dataset $\{(x_i, y_i)\}_{i=1}^n$, the model prediction $\hat{y}_i$ at iteration t is updated as follows:

$$\hat{y}_i^{(t)} = \hat{y}_i^{(t-1)} + \eta f_t(x_i), \quad (1)$$

where η is the learning rate, and $f_t(x_i)$ is the weak learner fitted to the residual errors from the previous iteration. The objective function to be minimized is given by:

$$\mathcal{L}^{(t)} = \sum_{i=1}^n L(y_i, \hat{y}_i^{(t)}) + \sum_{k=1}^t \Omega(f_k), \quad (2)$$

where $\Omega(f_k)$ is a regularization term that penalizes the complexity of the model, helping to prevent overfitting.

4.2 XGBoost Algorithm

XGBoost extends the gradient boosting framework with advanced features, including regularization, parallelization, and handling of missing values. The objective function in XGBoost can be written as:

$$\mathcal{L} = \sum_{i=1}^n l(y_i, \hat{y}_i^{(t-1)} + f_t(x_i)) + \Omega(f_t), \quad (3)$$

where l is a differentiable convex loss function that measures the difference between the prediction and the target. The regularization term $\Omega(f_t)$ is defined as:

$$\Omega(f_t) = \gamma T + \frac{1}{2} \lambda \sum_{j=1}^T w_j^2, \quad (4)$$

where T is the number of leaves in the tree, w_j is the weight of leaf j , γ is a regularization parameter for the number of leaves, and λ is a regularization parameter for the leaf weights.

In estimating ESSER spending on tutoring, we utilized the trained XGBoost model to predict the expenses based on features such as school enrollment, number of schools, program allocations, and categorical variables for states and locales. The dataset was cleaned to remove any rows with missing or zero values for tutoring expenses, and outliers were filtered using the Interquartile Range (IQR) method. Categorical variables were then encoded as dummy variables. After training the XGBoost model, we calculated the residuals, defined as the difference between the actual values y and the predicted values $\hat{y}$. The standard deviation of these residuals, σ_{res} , was used to establish an error margin for the predictions. For each prediction, the low estimate was calculated as $\hat{y} - \sigma_{\text{res}}$ and the high estimate as $\hat{y} + \sigma_{\text{res}}$, providing a range that accounts for the variability in the model's predictions.

4.3 Model Training and Evaluation

We utilized the XGBoost implementation from the `xgboost` Python library. The following steps outline the process:

1. **Data Preparation:** The dataset was cleaned by removing rows with missing or zero values for tutoring expenses. Categorical features were encoded using one-hot encoding.
2. **Train-Test Split:** The data was split into training and test sets using an 80-20 split.

3. **Model Training:** The XGBoost model was trained on the training set using 100 estimators. The hyperparameters were chosen based on standard recommendations and adjusted through cross-validation.
4. **Model Evaluation:** The model’s performance was evaluated on the test set using metrics such as Mean Absolute Error (MAE), Mean Squared Error (MSE), Root Mean Squared Error (RMSE), R^2 Score, Adjusted R^2 Score, and Mean Absolute Percentage Error (MAPE).

We employed an Optuna-based hyperparameter search within an XGBoost regression framework to estimate the share of ESSER funds allocated to high-dose tutoring (HDT). In this approach, we defined a scalar-valued objective function that performs five-fold cross-validation on the training set using the XGBRegressor estimator. Specifically, the objective function receives a set of hyperparameters from the trial, instantiates an XGBRegressor with those hyperparameters, and returns the negative mean of the cross-validated root mean squared errors (RMSE). The search space encompassed nine parameters: `n_estimators` (integer, 50–1000), `max_depth` (integer, 3–15), `learning_rate` (float, 10^{-4} –0.5, log-scale), `subsample` (float, 0.5–1.0), `colsample_bytree` (float, 0.5–1.0), `min_child_weight` (integer, 1–7), `reg_alpha` (float, 10^{-8} –1, log-scale), `reg_lambda` (float, 10^{-8} –1, log-scale), and `gamma` (float, 10^{-8} –1, log-scale). The hyperparameters were explored via Tree-structured Parzen Estimator (TPE), which adaptively models the density of favorable versus unfavorable hyperparameter regions. Each trial yielded a distinct hyperparameter combination, and the scoring function reported the mean of five-fold cross-validated RMSE values, ensuring robust estimation of generalization error. The study was executed for 250 trials, after which the optimized hyperparameter set was extracted via `study.best_params`. Lastly, we serialized the values of the objective function across all trials for subsequent inspection and reproducibility.

5 Results

The results are organized into three components. First, we assess the model’s internal fit and out-of-sample performance during training and validation. Second, we apply the trained model to the broader set of districts that mention tutoring but do not report expenditure amounts, in order to generate predicted spending estimates. Third, we analyze residuals from the fitted model to quantify uncertainty and construct an estimated range for total district-level allocations toward high-dosage tutoring.

5.1 Model Fitting and Training

Model fitting refers to the process by which the XGBoost algorithm iteratively minimizes prediction error on the training data by adjusting its internal parameters to capture statistical relationships between explanatory variables and the response. In this study, the model was trained on districts that explicitly reported tutoring expenditures, using features such as total ESSER allocation, student enrollment, number of schools, and urbanicity. The training process involved repeated gradient-based optimization to reduce the residual difference between predicted and observed expenditures, while the hyperparameter search ensured an optimal balance between model flexibility and generalization.

The hyperparameter optimization procedure, conducted through 250 Optuna trials, converged on a configuration characterized by moderate depth and conservative regularization. The final model used 457 estimators (`n_estimators`=457), a tree depth of 3 (`max_depth`=3), and a learning rate of 0.0108, promoting incremental updates that reduce the risk of overfitting. Subsampling rates for rows and features (`subsample`=0.7906; `colsample_bytree`=0.8462) introduced stochasticity to stabilize the model, while minimal regularization parameters (`reg_alpha`, `reg_lambda`, and `gamma` near 10^{-7}) indicated that the model achieved satisfactory performance without requiring heavy penalization of complexity.

The fitted model explained approximately 35.8% of the variance in reported tutoring expenditures ($R^2 = 0.3581$), with an adjusted $R^2 = 0.3361$ accounting for model complexity. While a portion of variance remains unexplained—reflecting the heterogeneity of district reporting and local fiscal practices—these results indicate that the model effectively captures systematic relationships among key district-level predictors of tutoring expenditure.

5.2 Application to Unreported Cases

Following training, the fitted model was applied to districts that referenced tutoring in their ESSER plans but did not disclose a dollar amount. In this application phase, the model’s parameters were held constant; no additional fitting or optimization was performed. Instead, the model produced predicted values based on the same input variables that were found to be predictive in the training data.

This application yielded an estimated total of approximately \$2.2 billion in district-level allocations for high-dosage tutoring. While this figure does not represent a definitive accounting total, it provides an empirically grounded estimate

derived from observed data and model-inferred relationships. Given the wide range of anecdotal estimates in circulation, spanning from roughly \$700 million to \$7.5 billion, this result offers a data-driven midpoint anchored in reproducible statistical inference.

5.3 Residual Analysis and Uncertainty Estimation

To evaluate the reliability of these estimates, we examined the distribution of residuals from the fitted model. In other words, we inspect the differences between observed and predicted tutoring expenditures among districts with known values. As shown in Figure 3, the residuals exhibit a narrow, symmetric distribution centered near zero, with most deviations falling between \$0 and \$20,000. This distribution suggests that the model’s predictive errors are both limited in magnitude and unbiased in direction.

While the majority of residuals cluster tightly around zero, several pronounced outliers are visible in the upper tail of the distribution. These high-magnitude residuals indicate a small number of districts where the model substantially underestimates actual expenditures. The skew toward positive residuals suggests that, in these instances, reported spending exceeded the level predicted by the model—potentially reflecting unique local budget priorities, multi-year tutoring contracts, or reporting anomalies. Although infrequent, these outliers underscore the heterogeneity of district-level spending behavior and warrant consideration when interpreting aggregate estimates.

We used the standard deviation of residuals (σ_{res}) as a measure of typical prediction uncertainty. By adding and subtracting this value from each model prediction, we derived upper and lower bounds for total predicted expenditures. This approach yields an estimated range of \$1.87 billion to \$3.85 billion in total district-level spending on high-dosage tutoring. The width of this range reflects uncertainty stemming primarily from data incompleteness and the inherent heterogeneity of district plans.

5.4 Interpretation

Taken together, these analyses reveal consistent evidence that districts prioritized high-dosage tutoring as a central strategy for academic recovery. Even under the most conservative assumptions, estimated allocations approach \$2 billion nationwide; at the higher end, they exceed \$3.8 billion. This finding positions high-dosage tutoring as a substantively and fiscally significant category of ESSER-funded interventions.

More, the modeling approach demonstrates that machine learning techniques can reconstruct credible expenditure patterns from incomplete administrative data. The model’s performance, coupled with the transparency of its validation, provides a replicable framework for analyzing large-scale education finance data where formal reporting mechanisms are incomplete or inconsistent. In doing so, it bridges an important methodological gap between policy analysis, inferential analysis, and machine learning in education research.

6 Scholarly Significance

This study advances our understanding of how federal recovery funds have been deployed to combat pandemic-related learning loss, offering the first systematic, data-driven estimate of district-level investments in high-dosage tutoring. Using a tuned XGBoost regression model on a uniquely constructed dataset of over 8,000 ESSER plans, we demonstrate how machine learning methods can illuminate opaque patterns of educational spending that elude conventional reporting systems. The model’s explanatory power and predictive precision provide an empirically grounded estimate of roughly \$2.2 billion in ESSER allocations to HDT, a figure that, while bounded by uncertainty, anchors the national conversation in evidence rather than conjecture.

The implications are twofold. First, this analysis demonstrates that a substantial share of relief funding was directed toward one of the most empirically validated interventions available. Methodologically, this work contributes to the applied machine learning literature by demonstrating how gradient boosting can be used for fiscal inference in under-structured administrative datasets. The framework generalizes to other educational or policy contexts where missingness and heterogeneity hinder econometric analysis. Together, these contributions suggest that even in fragmented policy environments, statistical learning techniques can restore a degree of transparency and accountability to public education finance.

Yet, the findings also expose the limits of current data infrastructure. The wide confidence range surrounding estimated tutoring expenditures underscores a fundamental problem: despite historic federal investment, no unified reporting system exists to track how funds were actually used. Without standardized, machine-readable documentation of district spending, future researchers and policymakers alike will continue to operate in partial darkness.

In sum, this work demonstrates both what is possible and what remains undone. By applying advanced analytic methods to an incomplete but vital record, we show how high-dosage tutoring has become a cornerstone of pandemic recovery efforts—and how much more could be learned, and accomplished, with transparent, interoperable data systems. The path forward lies in ensuring that the evidence we already have can be seen, shared, and acted upon.

References

- Megan Kuhfeld, James Soland, Beth Tarasawa, Angela Johnson, Erik Ruzek, and Jing Liu. Projecting the Potential Impact of COVID-19 School Closures on Academic Achievement. *Educational Researcher*, 49(8): 549–565, November 2020. ISSN 0013-189X, 1935-102X. doi:10.3102/0013189X20965918. URL <http://journals.sagepub.com/doi/10.3102/0013189X20965918>.
- Emma Dorn, Bryan Hancock, Jimmy Sarakatsannis, and Ellen Viruleg. COVID-19 and student learning in the United States: The hurt could last a lifetime. *McKinsey & Company*, 1:1–9, 2020. URL https://www.childrensinstitute.net/sites/default/files/documents/COVID-19-and-student-learning-in-the-United-States_FINAL.pdf.
- Andre Nickow, Philip Oreopoulos, and Vincent Quan. The Impressive Effects of Tutoring on PreK-12 Learning: A Systematic Review and Meta-Analysis of the Experimental Evidence, July 2020. URL <https://www.nber.org/papers/w27476>.
- Matthew A. Kraft, John A. List, Jeffrey A. Livingston, and Sally Sadoff. Online Tutoring by College Volunteers: Experimental Evidence from a Pilot Program. *AEA Papers and Proceedings*, 112:614–618, May 2022. ISSN 2574-0768. doi:10.1257/pandp.20221038. URL <https://www.aeaweb.org/articles?id=10.1257/pandp.20221038>.
- Tianqi Chen and Carlos Guestrin. XGBoost: A Scalable Tree Boosting System. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pages 785–794, San Francisco California USA, August 2016. ACM. ISBN 978-1-4503-4232-2. doi:10.1145/2939672.2939785. URL <https://dl.acm.org/doi/10.1145/2939672.2939785>.
- Dan Dewey, Erin Fahle, Thomas J. Kane, Sean F. Reardon, and Douglas O. Staiger. Federal Pandemic Relief and Academic Recovery. Technical report, Center for Education Policy Research, Harvard University, June 2024.
- Carly D. Robinson, Matthew A. Kraft, Susanna Loeb, and Beth E. Schueler. Accelerating Student Learning with High-Dosage Tutoring. EdResearch for Recovery Design Principles Series. *EdResearch for recovery project*, 2021. URL <https://eric.ed.gov/?id=ED613847>. Publisher: ERIC.
- Maria V. Carbonari, Daniel Dewey, Thomas J. Kane, Atsuko Muroga, Michael DeArmond, Elise Dizon-Ross, Dan Goldhaber, Emily Morton, Miles Davison, and Ayesha K. Hashim. The Impact and Implementation of Academic Interventions During COVID: Evidence from the Road to Recovery Project. 2024.
- Julia Gwynne and Elaine Allensworth. Mitigating the Impact of the COVID-19 Pandemic through Curriculum-Based Approaches to Learning Acceleration in Grades K-2 Literacy in Chicago. May 2023. doi:10.17605/OSF.IO/YNEFQ. Publisher: OSF.
- Robert Balfanz and Vaughan Byrnes. Increasing School Capacity to Meet Students’ Post-Pandemic Needs. *Everyone Graduates Center*, 2023.
- Jens Dietrichson, Martin Bøg, Trine Filges, and Anne-Marie Klint Jørgensen. Academic Interventions for Elementary and Middle School Students With Low Socioeconomic Status: A Systematic Review and Meta-Analysis. *Review of Educational Research*, 87(2):243–282, April 2017. ISSN 0034-6543, 1935-1046. doi:10.3102/0034654316687036. URL <http://journals.sagepub.com/doi/10.3102/0034654316687036>.
- Jonathan Guryan, Jens Ludwig, Monica P. Bhatt, Philip J. Cook, Jonathan M. V. Davis, Kenneth Dodge, George Farkas, Roland G. Fryer Jr., Susan Mayer, Harold Pollack, Laurence Steinberg, and Greg Stoddard. Not Too Late: Improving Academic Outcomes among Adolescents. *American Economic Review*, 113(3):738–765, March 2023. ISSN 0002-8282. doi:10.1257/aer.20210434. URL <https://www.aeaweb.org/articles?id=10.1257/aer.20210434>.
- Ahlam Almalawi, Ben Soh, Alice Li, and Halima Samra. Predictive models for educational purposes: A systematic review. *Big Data and Cognitive Computing*, 8(12):187, 2024. URL <https://www.mdpi.com/2504-2289/8/12/187>. Publisher: MDPI.
- Soukaina Hakkal and Ayoub Ait Lahcen. XGBoost to enhance learner performance prediction. *Computers and Education: Artificial Intelligence*, 7:100254, 2024. URL <https://www.sciencedirect.com/science/article/pii/S2666920X24000572>. Publisher: Elsevier.

- Pedro Carmona, Francisco Climent, and Alexandre Momparler. Predicting failure in the US banking sector: An extreme gradient boosting approach. *International Review of Economics & Finance*, 61:304–323, 2019. URL <https://www.sciencedirect.com/science/article/pii/S1059056017306950>. Publisher: Elsevier.
- Mark Eshwar Lokanan and Sana Ramzan. Predicting financial distress in TSX-listed firms using machine learning algorithms. *Frontiers in Artificial Intelligence*, 7:1466321, 2024. URL <https://www.frontiersin.org/journals/artificial-intelligence/articles/10.3389/frai.2024.1466321/full>. Publisher: Frontiers Media SA.

A Figures and Tables

Overlap of NCES IDs between Number Data and Word Data

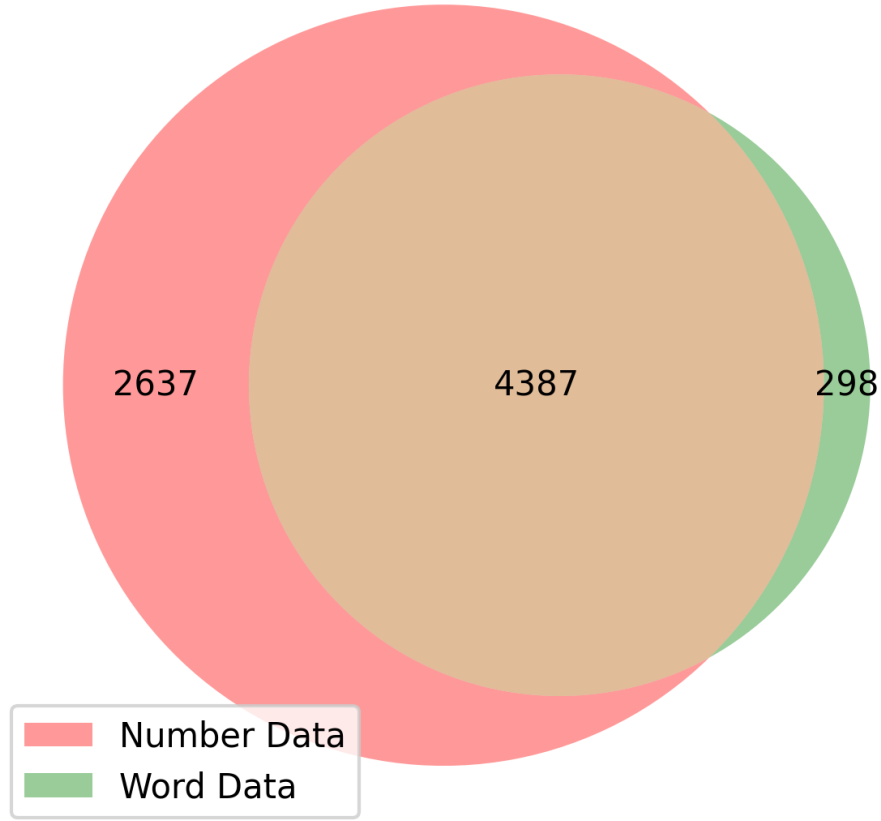


Figure 1: The overlap between the districts that have been manually annotated for budgetary items "Number Data," and the districts that have been passed through OCR and can be sorted by the presence/absence of tutoring "Word Data"

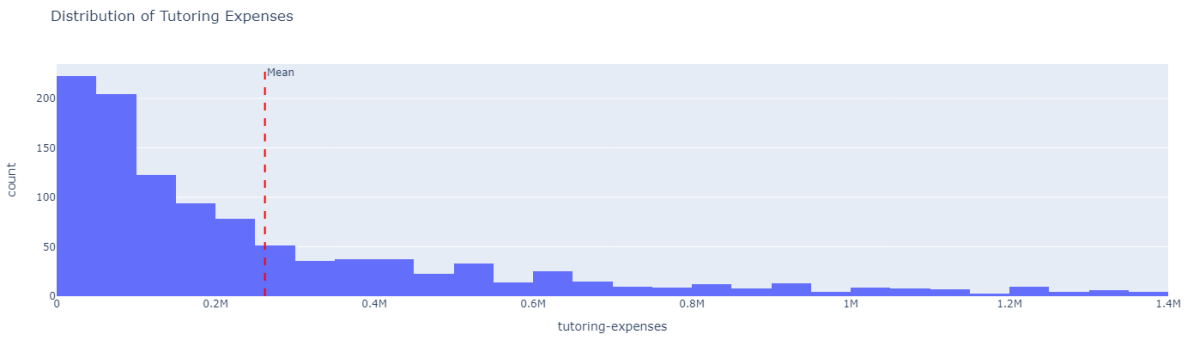


Figure 2: Distribution of tutoring expenses. Each bin is a \$50,000 increment.

Parameter	Min	Max
n_estimators	51	948
max_depth	3	15
learning_rate	1.06e-4	0.327
subsample	0.50	0.986
colsample_bytree	0.50	0.991
min_child_weight	1	7
reg_alpha	1.01e-8	0.919
reg_lambda	1.34e-8	0.796
gamma	1.31e-8	0.645

Table 1: Overall parameter search space and observed ranges. A total of 150 iterations were conducted.

Iter.	Value	n_estim.	max_depth	learn_rate	min_child_w.
15	1.601757e+6	51	6	0.0178	3
42	2.141770e+6	465	9	0.0806	4
64	1.801879e+6	708	3	1.07e-4	4
71	1.605657e+6	342	6	0.00072	4
99	2.121397e+6	775	15	0.00046	7

Table 2: Five illustrative iterations. We include one near the best (lowest) value, one near the worst (highest) value, and others that highlight extremes in hyperparameters such as very low learning_rate or large n_estimators.

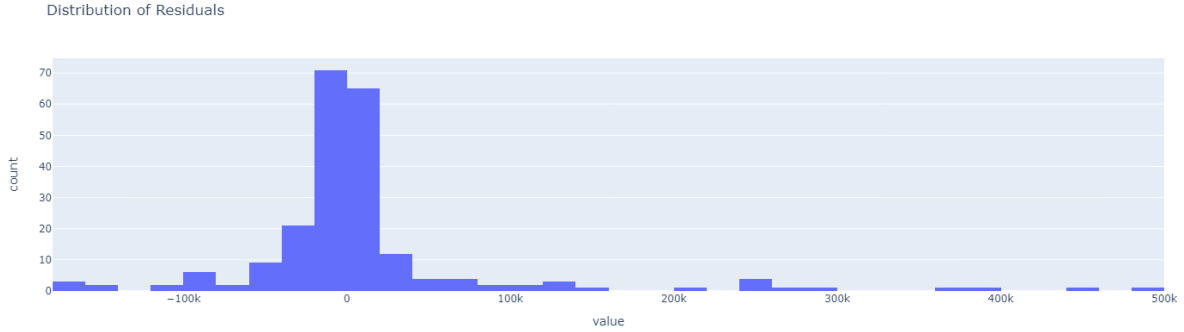


Figure 3: Distribution of residuals. Each bin is a \$20,000 increment.

Note: These iterations were chosen to cover the best and worst performance, plus cases of extreme parameter settings that illustrate the diversity of the search.