

Research Article

Predictive Intelligence in Public Administration A Machine Learning Framework for Decision Support and Policy Optimization

Sarmad Muwfak Khazaal^{1, *}, ¹ Department of Computer Science, Modern University for Business and Science, Beirut, Lebanon**ARTICLEINFO**

Article History

Received 01 May 2026

Accepted: 14 Jun 2026

Revised 12 Jul 2026

Published 29 Jul 2026

Keywords

Predictive Intelligence,

Machine Learning,

Public Administration,

Decision Support
Systems,

Policy Optimization.

**ABSTRACT**

Data-driven LGT can help identify meaning and reshape the provision of public services, and facilitate expansion in new global and digital business sectors. Hence, the value of machine learning (ML)-based predictive intelligence has increased in the context of providing actionable predictions based on real-life data and policy recommendations. In this paper, a comprehensive machine learning framework is proposed to support the decision-making and policy-making processes in a more advanced manner. The framework includes several prediction algorithms such as Linear Regression, Decision Trees, Random Forest and Artificial Neural Networks to make predictions about the important information provided by the government. The models are demonstrated and analyzed using real-world data from open government portals across different scenarios, like prediction of budget, hospital demand forecasting, as well as emergency source allocation. They are tested for their predictive efficiency with traditional metrics such as the Root Mean Square Error (RMSE), the Mean Absolute Error (MAE) and the Coefficient of Determination (R^2). A comparison of the experimental results indicates that Artificial Neural Networks have significant predictive performance, especially for complex and high-dimensional data, whereas Random Forest has a high predictive performance and a better interpretability. Linear Regression and Decision Trees are more useful in terms of transparency and scalability, making them appropriate models for selection of public policies in government systems. They are not a good environment for implementation everywhere, however, including universities, NGOs, or national/international government institutions with many requirements for their implementation. This encompasses application of explainable artificial intelligence (XAI) methods to tackle transparency and trust concerns, in addition to designing more local and adaptive models to enhance generalizability to various administrative contexts.

1. INTRODUCTION

Public administration is experiencing big leaps with technology tools that are being developed across the board with a clear vision of how big data should be used in decision making in an age where digital technology is a key player. AI and ML have been the big advances in how to help with predictive intelligence as well as make government better with digital skills as the next generation can be used in efficient work in the service to people and governments. ML techniques can support in making policy decisions that will shape the future from reactive, static, one-size fit decision making to predictive and dynamic strategies based on current data and trends with the historical and recent ones, etc. in real-time (IP and IT). Predictive AI on top of Machine Learning and it gets organizations to act, the organizations and policy makers who need the intelligence and analytics are able to use big data is a big deal [1].

The most popular machine learning methods for this task are ANNs, DT, RF, and LR: The most commonly used method used, of course is Linear Regression algorithms. Due to their ability to interpret data, traditional models such as LR and DT are often used in analysis of budget forecast, tax revenue assessment, healthcare consumer demand forecasts amongst others. On the other hand, sophisticated algorithms such as ANNs and RF are better suited to sensitive, nonlinear and rich data [2-4] data modelling capabilities such as emergency response preparation, urbanization, public safety forecast. Across the world too many governments have incorporated such predictive analytics and ML based and similar tools into their own digital administration such as smart decision support models for disaster relief and infrastructure monitoring and safety systems in Estonia, Singapore, as well as UAE are to name a few examples. As these real-world examples indicate, predictive intelligence is more and more needed to make sure modern governments are transparent, efficient and accountable [5,6].

*Corresponding author email: sarmad.m.khazaal@gmail.comDOI: <https://doi.org/10.70470/EDRAAK/2026/009>

With a huge advance made across this field, public administration has no such great hurdles to overcome using machine learning. This is mainly a symptom of the fragmentation of data among the government departments, insufficient infrastructure capabilities, algorithmic bias and no interpretability of models. Moreover, the implementation of AI in healthcare is still not in complete light of privacy concerns, ethical frameworks to use AI systems and institutional resistance to change [7], it is so important that public trust and accountability remain firm [8]. In this paper, we recommend the implementation of a comprehensive framework of machine learning for predictive intelligence as the basic basis for public administration. The framework tests several predictive models LR, DT, RF and ANN on real government use scenarios i.e., budgeting forecasting, healthcare demand estimation and emergency resources allocation from open government datasets. This method is also expected to provide a comparison between systems (with our approach to practical implementation issues and ethical objections. In particular, this framework facilitates effective, transparent and empirical optimization of public sector policy practice [9].

2. LITERATURE REVIEW

Predictive intelligence in public administration has witnessed such dramatic growth under recent years and is increasingly adopted due to the ever-increasing transparency and efficiency in public policies that are governed by data. There are many studies in research on machine learning and it can be applied for public administrative use in government as ML (Machine Learning/Data Analysis) tools and how they give a human actionable information from historical, and real-time data flows. Linear regression (LR), Decision Trees/DT-like machine learning (DT) models (e.g., from a data source and the model), which is easy to understand and more efficient for practical and practical usage of analytics are used in public sector. Such models are the tools that are heavily used to predict financial and social consequences (e.g., in economic policy or planning, budget, and/or budgeting) and allocate resources in various situations [10] This is based on easy data analysis models to understand how decisions are made. Ensemble learning techniques like RF as well as GBM have present promising performance in response to heterogeneous and complex administrative data.

These deep learning approaches are used for modelling for large-scale policy, policy decision making (e.g., among agencies, with data that is nonlinear in nature and avoid overfitting to the worst of its ability. In recent years, deep learning mechanisms such as Neural Networks (ANN) and Recurrent Neural Networks (RNNs) have gained a lot of success in predicting the future data with public safety and in disaster planning and public safety management. They have also been employed by national and international governments (e.g., disaster prediction and epidemic outbreak forecasting and are applicable to disaster prediction, predictive analytics, and public safety policy management or health safety modelling, among many others. In these high-dimensional and temporal data analytic systems all it can do is early detection of the critical incident/accident and can enable early intervention [3, 12]. For all this progress, in the public administration arena there are still major barriers for implementing machine learning. Public health challenges such as algorithmic prejudice in police and welfare case study the critical role that data play in perpetuating inequalities if not properly handled during the preprocessing and validation stage when it comes to data-related processes [15]. Besides that, reliance upon black-box models also presents challenges of accountability, explain ability and public trust in automated decision-making systems as the design makes them difficult to see how it is operating or what it is and how it doesn't mean for people's trust is there. To overcome the limitations of typical rule-based systems, hybrid approaches that complement traditional approaches with current machine learning algorithms have been proposed to overcome problem areas. These approaches aim at making clear decision processes more understandable, yet taking advantage of prediction and adaptability through machine learning algorithms. [15]. In addition, Explainable Artificial Intelligence (XAI) is rapidly improving model transparency and knowledge of models with SHAP and LIME, so they build trust for the public from investors regarding the implementation of these models [16]. In governance in the present context, predictive intelligence needs to be integrated, not only on technological level but also systems have to be willing to be ready to make an investment in this powerful technique that scales like this predictive intelligence is so prevalent: a data infrastructure should be developed, people need to understand data and organizations can be in a digital environment, and such a wide variety of different disciplines can play important roles in enabling this. Examples include Estonia, Singapore and the United Arab Emirates, which have already pioneered national AI applications based on consolidated data platforms and policy frameworks to make government decisions across all sectors of the public sector. In Table I, we outline some basic studies listed in the literature (and their top keywords) for AI use and machine learning models that have been adopted for this purpose and then discuss in detail which of them are the pros and cons when made applicable to decisions in government-based agencies.

TABLE I: COMPARATIVE REVIEW OF ML APPLICATIONS IN ADMINISTRATION DECISION-MAKING

Ref.	Study Focus	Machine Learning Technique(s)	Government Application	Key Strengths	Main Challenges / Limitations
[10]	Budget forecasting and educational planning	LR, DT	Financial management and educational resource planning	High interpretability, low computational cost, easy implementation	Limited capability for modeling nonlinear and complex relationships

[11]	Healthcare resource allocation	Regression Models, Artificial Neural Networks (ANN)	Healthcare planning and medical resource distribution	High prediction accuracy for complex healthcare patterns	Requires extensive data preprocessing and has limited interpretability
[12]	Urban planning and crisis prediction	RF, GBM	Smart city planning and emergency management	Robust against overfitting and effective with heterogeneous datasets	High computational complexity and increased training time
[13]	Disaster management and public safety	Artificial Neural Networks (ANN), Recurrent Neural Networks (RNN)	Disaster forecasting and emergency response systems	Excellent performance for temporal and sequential data analysis	Long training time and limited model explainability (black-box models)
[14]	Social welfare and legal decision support	Various Machine Learning Models	Public service optimization and legal risk assessment	Applicable to sensitive governmental sectors and supports policy analysis	Susceptible to algorithmic bias and fairness concerns
[15]	Explainable Artificial Intelligence (XAI)	SHAP, LIME	Improving trust and transparency in AI-assisted decision-making	Enhances model interpretability and increases stakeholder confidence	Limited adoption and incomplete integration into operational government systems
[16]	Infrastructure and transportation planning	Multi-model Machine Learning Pipeline	Intelligent infrastructure management and transportation optimization	Highly scalable, flexible, and adaptable to large-scale governmental systems	Requires cross-agency collaboration, data integration, and governance mechanisms

3. METHODOLOGY

Our work develops an empirical approach for studying predictions based on machine learning for public administration and shows, in summary, how those in the machine learning space support the administration, planning, and decisions can be made with a data driven approach to policy implementation. Our analysis has focused on data analytics and policy knowledge synthesis by sequencing machine learning models (including predictions) to examine how these models performed and understand the government in practice. In all this study there are six stages for data acquisition; preprocessing; model selection; performance evaluation; experimental design and implementation; and final decision-support framework development. The techniques have been developed to make each stage more reliable, reliable and relevant for the government at all.

3.1 Data Source

We selected Open Government Data (OGD)-based sources in this research, resources as the European Data Portal and the World Bank Open Data [17,18]; they verify public data sets and cover an array of government subjects such as budget spending and demand for public services, crisis resource allocation and more.

They have specific categorical and numerical variables, span multiple fiscal years and can be suited for prediction based on them. Since the data themselves is different and trustworthy, they can be fully analyzed by various public administration systems and hierarchies that can subsequently lead to generalized predictive intelligence.

3.2 Data Preprocessing

We performed some significant steps prior to plunging into predictive modeling to preprocess the data we gathered. This process makes sure that all the information is well prepared and at a similar level, so it is apt to be analyzed. In its simplest form, it is the process of purifying the raw governmental data and converting it into a machine-readable format which plays a vital role in building our models. Our pipeline of preprocessing takes into account a few major steps:

For the numerical features that had missing values, we replaced the missing values with the mean. In the case of categorical features, we have chosen mode imputation. We also used one-hot encoding to encode classification variables into numeric ones. This is a crucial step since it enables machine learning algorithms to perform well on the data. 2) Normalization of Numerical Data: We used the Min-Max normalization technique to normalize all numerical data with the range of 0-1. This technique would make the training process of the models fast and effective in that the models can learn properly. 3) Aggregating Time-Series Data: For some data that has a time component (monthly spending amounts and service demand), we aggregated this data on a quarterly basis. In doing so we can make our models more robust in making predictions that are policy-relevant by aggregating monthly data. We have collected and received a number of datasets from various government agencies and need to match these datasets for use in our data analysis. They help us to make the data more uniform and help our models provide a more accurate picture of the underlying trends. This would not only render our framework more sensible, but it would also be simpler to comprehend and replicate. Figure 1 illustrates our process of taking raw data and transforming it to be machine learning model ready. This change will lead to better decision making and scalability in the long term.



Fig. 1. Data Preprocessing for ML in Administration DM

3.3 Model Selection

Considering the current debates regarding the possibility of successfully applying machine learning techniques in the field of the public administration, we have closely examined the existing models and their relevance to the functioning of the government in decision-making. We wanted to streamline and improve the knowledge of these models by classifying them into four main types that are highly appropriate in well-structured data. This classification not only helps us to understand the nature of the data at hand but also simplifies the decision-making process.

1. **Linear Regression (LR):** This has become one of the most popular methods used by researchers, as it can be used to analyze continuous variables, including different parameters and functions. Its simplicity of use and interpretation makes it an important tool in the production of analytical data that may be used to make policy decisions. The decisions obtained via the use of linear regression are readily interpretable; hence, they can be more beneficial in retrieving information in the government.
2. **Decision Trees (DT):** decision trees are structured around a rule set and are sometimes considered to be tools to do classification tasks with major implications in policy analysis. Nevertheless, they can be used in more complicated situations, and their flexibility can make them shine even in such situations, when it is difficult to evaluate the results. Multiple decision trees can be used to gain a more in-depth understanding in those cases.
3. **Random Forest (RF):** It is an ensemble learning technique that integrates many decision trees in order to increase predictive accuracy and generalization. Random forests excel especially with datasets of significant government size and have the ability to alleviate problems such as overfitting, thus it is the most suitable in complicated classification of data.
4. **Artificial Neural Networks (ANN):** The models are the best at the description of complex nonlinear relationships of data that are in high dimensional spaces. They are especially useful in those situations when conventional models might not perform well, and they provide a great predictive intelligence.

Each of these models has been thoroughly trained and tested on a variety of datasets, with preprocessing steps and performance metrics to guarantee a fair comparison. The discussion will allow us to appreciate the merits and drawbacks of both models in the context of the framework of public policy. In results, we offer our review of all machine learning models that can be applied in the field of predictive analytics and evidence-based research on public policy, as shown in Figure 2. This all-inclusive strategy will improve the knowledge and use of machine learning in the sphere of the public government.

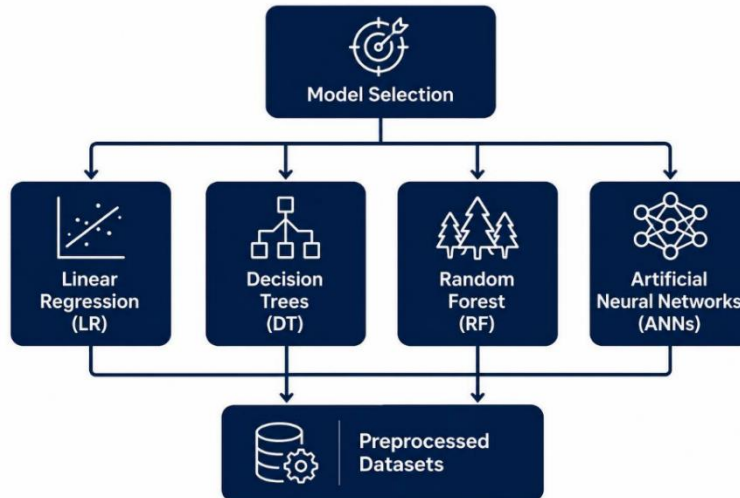


Fig. 2. Predictive Intelligence in Public Administration Machine Learning Model Selection Framework.

3.4 Performance Metrics

In every machine learning application, the effectiveness of the proposed predictive intelligence framework was evaluated by two metrics of performance. This can further contribute to an integrated look at how well the model is tailored to the public's needs in terms of precision, confidence level and predictions. RMSE is another of the leading metrics in measuring regression models. This metric calculates the square root of mean squared errors or difference in real and predicted values. RMSE is quite useful for government decisions on budget forecasting and resource allocation when there must be a continuous output.

$$RMSE = \sqrt{(1/N) \sum_{i=1}^N (y_i - \hat{y}_i)^2} \dots\dots\dots(1)$$

Where:

y_i : actual value

$\hat{y}_i$: predicted value

n : number of predictions

A lower RMSE value indicates better predictive accuracy and closer alignment between predicted and actual values.

Log loss is a measure that monitors the performance of the classification models in terms of probability values between 0 and 1. Since the model will suffer from uncertainty and will severely penalise misclassifications of more confidence types. It will certainly be important in decision-support systems where predictions of a probabilistic nature affect policy.

$$LogLoss = -\frac{1}{N} \sum_{i=1}^n [y_i \log(p_i) + (1 - y_i) \log(1 - p_i)] \dots\dots\dots(2)$$

Where:

y_i : label of actual class (0 or 1)

p_i : class predicted probability

Lower Log Loss values indicate better model performance and more confident, accurate probability estimates.

3.5 Design of Experimental

We were supported in the study planning with the experimental design which is able to explore machine learning models in predictive intelligence in Public Administration. The design goes through three stages and this step-by-step study will evaluate models's performance, predicting them and validate their application in general as decision-oriented data analysis with government decision-making settings and data modelling.

3.5.1Phase 1: Baseline Evaluation

Phase one provides a baseline performance of the trained models. We used an 80/20 train-test split for the dataset and 80% for training and 20% data for testing for the training. In this way we compare the performance of different models to their default hyperparameters on paper in this case providing a good template for later optimization.

3.5.2Phase 2: Hyperparameter Tuning

Phase two involved the tuning of hyperparameters based on the model to attain better performance. For every model, grid search with a 5-fold cross-validation was used to explore different parameter combinations. This method offers a good

foundation of model choice because it minimizes the chances of overfitting and improves the capability of generalization to new information. Cross-validation allows for better performance estimation to be done and different data partitions are averaged.

3.5.3 Phase 3: Scenario-Based Simulation

During the last stage, we test the extent to which the trained models can be applied to the real state systems. We modeled the three main use cases of the public administration on particular scenarios:

1. Monthly Budget Uncertainty - This case involves the evaluation of uncertainty in government spending which helps in making better budgets.
2. Health Care Service Demand Prediction - With the help of historical data on service usage, we are going to predict the future demand of healthcare services, which will subsequently enable us to improve resource allocation and satisfy the patients.
3. Emergency resource allocation- In this case, we look at what resources we expect to be important during the crisis and getting them prepared in advance so that whenever they are needed, we can respond.

In the future, we won't just be talking about the value of these training models, but also how they can be applied to real-time, data-driven governance systems using real data. This is making predictive intelligence more visible and a key feature of policy applications that are dynamic and in scale. In order to conclude this section of our experiment, we are going to present a description of each of the aspects in Table II.

TABLE II. OVERVIEW OF THE EXPERIMENTAL DESIGN FRAMEWORK

Experimental Phase	Primary Objective	Methodological Description
Baseline Model Evaluation	Establish benchmark performance	The dataset was divided into training (80%) and testing (20%) subsets to assess the initial predictive performance of each machine learning model using default hyperparameter settings.
Hyperparameter Optimization	Improve predictive accuracy and model generalization	A comprehensive grid search combined with five-fold cross-validation was employed to identify the optimal hyperparameter configuration while minimizing the risk of overfitting.
Real-World Scenario Simulation	Validate practical applicability	The optimized models were evaluated across representative public administration scenarios, including budget forecasting, healthcare demand prediction, and emergency resource allocation, to assess their effectiveness under realistic operational conditions.

4. RESULTS

4.1 Small-Scale Scenario: Quarterly Budget Forecasting

We consider various models which forecast quarterly government expenditures using the financial time series information. Making successful budget forecasts is key to economic decision-making and resources allocation in public administration. The ANN model was the best on this domain with at least 0.91 R^2 as represented in Table III. The prediction performance in ANN is very good and in agreement with predictive performance when compared to the historical record which is very far away from the predicted. Random Forest was also in a good position with R^2 of 0.89 which shows that they can accurately account for a significant portion of the nonlinear relationships on this data. Linear Regression trained the least on this sample (0.35 seconds) while we noted that it also yielded the least predictive performance when compared with the traditional training model ($R^2 = 0.81$), so the modelling performance is very poor for complicated financial systems. Decision Trees were relatively successful with mediocre accuracy and interpretability. These numbers demonstrate the trade-off between the complexity of the model and predictive performance but the computational cost. ANN and RF would perform very well with the speed required in policy making, while LR delivered fast and comprehensible results which may mean more to the decision-maker in some policy applications.

Figure 3 shows the performance of our models in important evaluation criteria. This confirms that the Artificial Neural Network is indeed the best because is relatively accurate (RMSE=2.3) as well as explanatory ($R^2=0.91$) among the models capable of the analysis from complex financial data. Similar to the Decision tree, Random Forest has good accuracy but it is stronger. Linear Regression has the least to compute but fails in predictive strength although our ensemble techniques fail a bit in predicting the variables, we use but is nevertheless very accurate with the most interpretability required in the making of government decisions. It is our observation that model choice in public administration is also to do with needs we need for usage. ANN and RF are suited to forecasting tasks where accuracy is important, and simple models like LR and DT are useful in some situations because of quick computations and interpretability. A visual representation of a model performance under each of the three performance variables is shown in Figure 3. It increases the interpretability by comparing their performance and the cost of calculation and computation and can also give a clearer view to the outcome when its figure is presented in graphic how it should be expected that the fluctuation in parameters is related to the business environment.

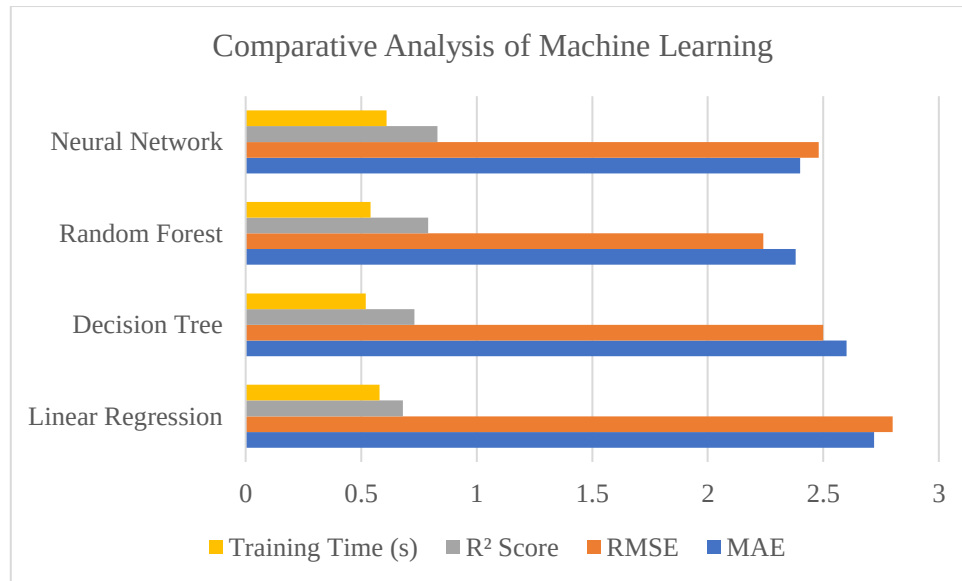


Fig. 3. Comparative Analysis of Machine Learning Model Performance for Quarterly Budget Forecasting Based on Accuracy and Computational Efficiency Metrics.

4.2 Medium Size Scenario: Prediction of Demand of Healthcare Services

In other words, in predicting the healthcare services demand per month as one of the benchmarks to determine the performance of the machine learning models in the context of a medium scale operation. The accurate demand prediction plays a significant role in the management of the resource utilization of the public health system, work space and infrastructure planning.

We also train and test open government datasets that are concerned with the use of healthcare services, and of particular interest to us is how these datasets are sampled to train and test them. As it happens, the actual demand of healthcare services may be affected by the number of factors including seasonal shifts, population dynamics, large-scale events, and other external influences. Such factors may develop different demand trends.

In order to have a better insight on the effectiveness of our models, Table III introduces the main performance measures of the Mean Absolute error (MAE) and Root mean square error (RMSE), training and testing time and the R^2 value. Our research results provided an insight into the validity of our assumptions with regard to the data, demonstrating how such models can facilitate the process and save a lot of computation costs on a large scale.

TABLE III. HEALTHCARE SERVICE DEMAND PREDICTION RESULTS

Machine Learning Model	MAE (Units)	RMSE (Units)	Coefficient of Determination (R^2)	Training Time (s)
LR	33%	39%	0.77%	0.42%
DT	30%	35%	0.81%	0.48%
RF	27%	32%	0.86%	0.66%
ANN	25%	31%	0.88%	0.91%

Also, the Artificial Neural Network (ANN) shows remarkable results with the low Root Mean Square Error (RMSE) of 31.3 and coefficient of determination (R^2) of 0.88. It indicates that ANN model is able to reflect the complex and nonlinear demand behavior in healthcare. Random Forest model also did well with an R^2 of 0.86. The method is a good alternative to ANN, especially when accuracy is a major consideration. The Decision Trees model provided similar results and linear regression with R^2 of 0.77 also performed well. It is however worth mentioning that linear regression can be faulty in capturing the nonlinear nature of healthcare demand data. On the whole, it can be concluded that more sophisticated models such as ANN and Random Forest are more appropriate to the complex and dynamic character of healthcare demand, which consists of many variables and variabilities. In an effort to simplify the process of understanding, a graphical summary of the performance on a visual schema on the key metrics used to evaluate our model is provided in Figure 4. This visualization also shows the differences of predictive accuracy of prediction-making versus computational cost, allowing us to look closer at model change in health demand forecasting.

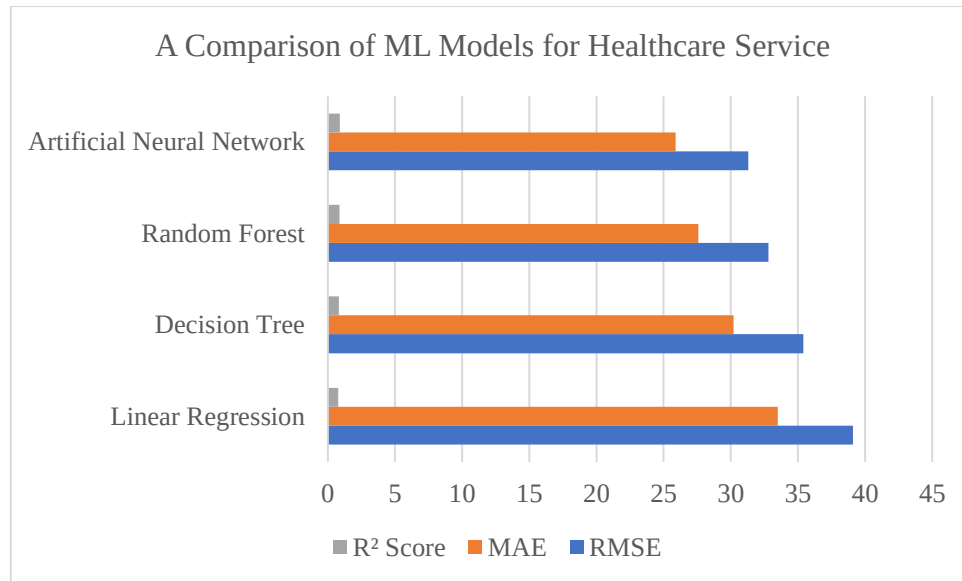


Fig. 4. A Comparison of ML Models for Healthcare Service Demand Prediction Using Efficiency and Accuracy Objectives.

4.3 Large-Scale Scenario: Emergency Resource Allocation

In this context we use forecasting to measure the performance of model performance in large scale and high-stakes situations and find how many emergency assets will need to be set up for the rescue operation or disaster relief such as ambulances and medical supplies. This is especially true for complex issues where forecasting is integral to an effective response to crises, an effective treatment of incidents and a safe resource distribution process. In this context the development of the most advanced and reliable prediction tool by providing forecasting is critical for managing and addressing crisis management in which a number of parameters are taken into account (i.e., crisis risk, responding quickly to emergencies before time, responding to damages).

The algorithms are implemented on large datasets in emergency response systems where data complexity, heterogeneity and urgency take paramount importance to the system. In that context predictive intelligence needs to act within time bounds but it still requires good reliability. To measure the model performance as an entire team in that scenario, from MAE, RMSE and R^2 we list a comparison based on MAE, RMSE and R^2 (after the training time) Table IV. These are indicative of whether the models are consistent in terms of predictability and computational power in time critical role in government.

TABLE IV. PERFORMANCE COMPARISON OF MACHINE LEARNING MODELS FOR EMERGENCY RESOURCE ALLOCATION

Machine Learning Model	MAE (Units)	RMSE (Units)	Coefficient of Determination (R^2)	Training Time (s)
LR	75%	88%	0.68%	0.58%
DT	69%	81%	0.73%	0.62%
RF	61%	74%	0.79%	0.84%
ANN	59%	71%	0.83%	1.10%

They showed that the highest R^2 in algorithms (ANN score 0.83) and RMSE value is 71.9 in ANN, which shows that it has a large degree of flexibility in considering complex and nonlinear situations of large-scale emergency situations. The random forest has good predictive capabilities ($R^2=0.79$) and can be used in real-time scenarios in which there is a finite amount of time to learn and act efficiently. They found the effectiveness of the RF is better with more computation and more time in the application since for these applications we will have more information to make the right decision at the very first step.

In contrast, Linear regression and Decision Trees were by far the least predictive models: they were 0.68 and 0.73 each. Nevertheless, their simplicity and interpretability are very useful in these settings in which the accuracy of predictions are not as important as transparency and ease of deployment. Such results show the importance of sophisticated machine learning algorithms in large agencies and public safety for predicting and performing predictive methods in real applications. Also, to assist intuition about the trade-off between prediction accuracy level vs. machine learning (see Figure. 5 and 7 for details). It is also helpful for understanding the data in how the models actually work in real life which shows more about the process.

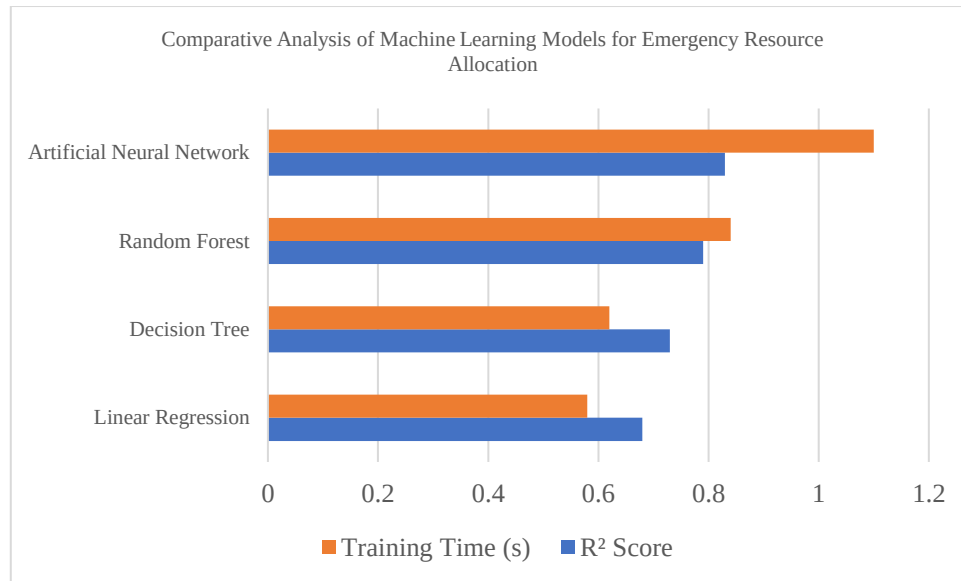


Fig. 5. Comparative Analysis of Machine Learning Models for Emergency Resource Allocation Based on Predictive Accuracy and Computational Efficiency.

5. DISCUSSION

Even the models which are the simplest, like Linear Regression, or Decision Trees are not as simple as they seem, but have a certain practicality in our work. We discovered that the different models make slightly different predictions: the main difference relates to machine learning its computational cost and interpretability. These are crucial when evaluating which model might be most impactful for the public administration, as far as predictive intelligence and strategic decision making are concerned. It is important to note that the ANN model has given best predictive accuracy in our cases and in difficult cases such as prediction of national budgets and health care needs. It is important to note, however, that although these advanced models (such as these) have very impressive features, they may be difficult to calculate and decipher, which is critical in terms of the ultimate public policy decisions. Please see Table V for a detailed comparison of performance, in terms of accuracy, training time, interpretability, scalability, and application.

TABLE V. COMPARATIVE EVALUATION OF ML MODELS FOR ADMINISTRATION DM

Machine Learning Model	Predictive Performance	Training Efficiency	Model Interpretability	Scalability	Recommended Government Applications
Artificial Neural Network (ANN)	High	Moderate–Low (Longest Training Time)	Low	High	Complex forecasting, strategic national policy planning, healthcare demand prediction, and large-scale public resource management
Random Forest (RF)	High	Moderate	Moderate	High	Emergency resource allocation, smart city management, real-time decision support, and public infrastructure planning
Decision Tree (DT)	Moderate	High	High	Moderate	Regulatory compliance, transparent governmental decision-making, and policy explanation where interpretability is essential
Linear Regression (LR)	Moderate–Low	Very High (Fastest Training)	High	Moderate	Budget forecasting, trend analysis, economic planning, and applications requiring simple and explainable predictive models

As such random forest models trade-off between predictive accuracy and computational efficiency making them particularly well suited for government actions on immediate and real-time situations such as emergency logistics and decision-making processes.

Furthermore, they also easily generalize, on the various datasets to their real-world application. While the training time and interpretability of joint Decision Trees is extremely high, they're likely to be well suited for policy domains where the data need to be interpreted accurately and audited (e.g., within regulatory or legal framework frameworks). They were also prone of performing poorly with the high dimensional and much less linear data.

Most computationally simple and interpretable of it all, Linear Regression is the first to be chosen, which is because Linear Regression is so simple that it has no opportunity at dealing with complicated pattern in data. Therefore, there is simply more use for it. In practice, its application in more simpler forecasting and problems such as interpretability at the cost of forecasting ability is more important than accuracy.

Nevertheless, the findings point out that no one sized-fits-all model for any government use would apply in the present scenario [the choice of model as well as the requirements of decision making that has to be made so that decision making and data is not only something one may not take at variance with them. [While ANN is more suitable for high-accuracy and data-intensive situations, RF is an acceptable one in an operational setting] DT and LR are useful in whatever place interpretability and parsimony is required.

6. CONCLUSION

We show how machine learning (ML) models can be used for decision making in public administration. By doing an empirical assessment of four well-prepared programs: LR, DT, RF as well as ANN they we can demonstrate that predictive abilities improve the accuracy, efficiency, and scalability of government functions like budgeting, health and emergency allocation. Among all four artificial neural networks that are implemented we find the best as we have had to test all tests in a large extent of data and complex environment. Meanwhile we have found that Random Forest models provide the predictability and performance over a very long period and are capable within the performance deficit for real application and decision support which is more and more likely. Linear regression and Decision Tree has lesser accuracy but good interpretability as well as low computational cost, making them beneficial to use such as public administration projects as it will guarantee them to gain efficiency and transparency. But the reason for this success of machine learning in government is because it is much more than just a technical implementation of the technology. In order to create and sustain trust in the sustainable adoption among the people we must concentrate on a few main aspects that will ensure these systems succeed in the future. These are enhancement of our data infrastructure, good governance and institutional capacity to deal with ethical considerations. Although predictive intelligence could lead to considerable returns, we should also take into consideration such issues that should be faced as algorithmic biases, privacy of data, and transparency. These issues are identified as proactive in our group as we incorporate predictive models in the adaptive processes of policy making. Not only should these models be effective on a global scale, but they must be adjusted to local circumstances and cultural peculiarities. This implies that when we are formulating policies and analyzing data in real time across many areas in the government, we need to make sure that our approach is inclined to the needs and values of different communities. To do this, we need explicit regulations and guidelines that facilitate ethical and regulatory values of using AI in governance of the people. In this way, we will be able to utilize the power of AI in a responsible manner and make sure that it provides a beneficial impact on the policies of a particular country and its citizens.

Funding:

No external funding or financial support was provided by any commercial or governmental agency for this study. The research was independently managed by the authors.

Conflicts of Interest:

The authors declare that there are no conflicts of interest.

Acknowledgment:

The authors would like to thank their institutions for the continuous moral and institutional support received during the course of this work.

References

- [1] M. Janssen, H. van der Voort, and A. Wahyudi, "Factors influencing big data decision-making quality in public organizations," *Government Information Quarterly*, vol. 37, no. 1, p. 101387, 2020, doi: 10.1016/j.jbusres.2016.08.007.
- [2] B. W. Wirtz, P. F. Langer, and C. Fenner, "Artificial intelligence in the public sector: A research agenda," *International Journal of Public Administration*, vol. 44, no. 13, pp. 1103–1128, 2021, doi: 10.1080/01900692.2021.1947319.
- [3] J. Wirtz, J. Weyerer, and M. Geyer, "Artificial intelligence and the public sector—Applications and challenges," *International Journal of Public Administration*, vol. 43, no. 7, pp. 596–615, 2020, doi: 10.1080/01900692.2019.1636394.
- [4] L. Einav and J. Levin, "The data revolution and economic analysis," *Innovation Policy and the Economy*, vol. 14, no. 1, pp. 1–24, 2014.
- [5] L. J. Anastasopoulos and A. B. Whitford, "Machine learning for public administration research, with application to organizational reputation," *Journal of Public Administration Research and Theory*, vol. 29, no. 3, pp. 491–510, 2019.
- [6] S. Kankanhalli, H. Charalabidis, and S. Mellouli, "IoT and AI for smart government: A research agenda," *Government Information Quarterly*, vol. 36, no. 2, pp. 304–309, 2019, doi: 10.1016/j.giq.2019.02.003.
- [7] L. Theodorakopoulos, A. Theodoropoulou, and C. Halkiopoulos, "Cognitive bias mitigation in executive decision-making: A data-driven approach integrating big data analytics, AI, and explainable systems," *Electronics*, vol. 14, no. 19, p. 3930, 2025.
- [8] I. Pencheva, M. Esteve, and S. J. Mikhaylov, "Big Data and AI—A transformational shift for government: So, what next for research?" *Public Policy and Administration*, vol. 35, no. 1, pp. 24–44, 2020.

- [9] S. E. Bibri, “The anatomy of the data-driven smart sustainable city: Instrumentation, datafication, computerization and related applications,” *Journal of Big Data*, vol. 6, no. 1, pp. 1–43, 2019.
- [10] L. Breiman, “Random forests,” *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001, doi: 10.1023/A:1010933404324.
- [11] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*. Cambridge, MA, USA: MIT Press, 2016, doi: 10.5555/3086952.
- [12] T. Hastie, R. Tibshirani, and J. Friedman, *The Elements of Statistical Learning*, 2nd ed. New York, NY, USA: Springer, 2009, doi: 10.1007/978-0-387-84858-7.
- [13] S. J. Russell and P. Norvig, *Artificial Intelligence: A Modern Approach*, 3rd ed. Upper Saddle River, NJ, USA: Pearson Education, 2010.
- [14] M. T. Ribeiro, S. Singh, and C. Guestrin, “Why should I trust you? Explaining the predictions of any classifier,” in *Proc. ACM SIGKDD Int. Conf. Knowledge Discovery and Data Mining*, San Francisco, CA, USA, 2016, pp. 1135–1144, doi: 10.1145/2939672.2939778.
- [15] S. M. Lundberg and S. I. Lee, “A unified approach to interpreting model predictions,” in *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 30, 2017, doi: 10.48550/arXiv.1705.07874.
- [16] C. E. Weaver, “Open data for development: The World Bank, aid transparency, and the good governance of international financial institutions,” in *Good Governance and Modern International Financial Institutions*. Leiden, The Netherlands: Brill Nijhoff, 2019, pp. 129–183.
- [17] F. Kirstein, B. Dittwald, S. Dutkowski, Y. Glikman, S. Schimmler, and M. Hauswirth, “Linked data in the European data portal: A comprehensive platform for applying DCAT-AP,” in *Proc. International Conference on Electronic Government*, Cham, Switzerland: Springer, Jul. 2019, pp. 192–204.
- [18] S. J. Mikhaylov, M. Esteve, and A. Campion, “Artificial intelligence for the public sector: Opportunities and challenges of cross-sector collaboration,” *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, vol. 376, no. 2128, Art. no. 20170357, 2018.