

TECHNICAL REPORT

A Comparative Study on High-Performance Carbon Footprint Prediction Using Interpretable Machine Learning

Khayla Naura Ulya Luqyana
Universitas Indonesia
Depok, Indonesia

Muhammad Aldo Fahrezy
Universitas Indonesia
Depok, Indonesia

Naufal Zafran Fadil
Universitas Indonesia
Depok, Indonesia

I. INTRODUCTION

Climate change and pollution has been a prevalent problem in humanity's ever growing technology and consumption, CO_2 being one of the primary culprit in many environmental issues. In the last 50 years, CO_2 emission has sextupled, reaching 37.79 billion tons in 2023 [5]. The United Nations even touted 'Responsible Production and Consumption' and 'Climate Action' as humanity's goal to achieve sustainable means of thriving and development [4]. [3] concluded consumers directly contribute to 50% of the greenhouse gas emissions generated. [8] suggested government and corporations to focus on developing environmentally freindly choices as the easy and favorable alternatives. Most studies on this topic covers implementation and discussion exclusively on specific sectors, rarely do so on an individual level [1]. Thus, our paper focuses on shaping more responsible consumer behavior's contribution to CO_2 emissions. We developed various models to capture the relation of each behavior towards CO_2 byproduct volume, as each behaviour can be attributed with direct and indirect contribution to CO_2 emission volume of varying degrees throughout the behavior or the related product's life cycle (e.g. albiet carbon-free, the production of electric vehicles (EV) still incorporates many CO_2 -intensive techniques). Therefore, we hope that the models we develop may capture linear and non-linear correlation of an individual's behavior with respect to carbon emission. We also look onto the model's computational efficiency and interpretability, as the model must have high throughput and performance in edge-device or low resource settings and provide actionable insights to the user, government, and corporations on combatting climate change respectively. Such model criteria allows us to uncover valuable insights among vast data on related issues. Our result may also shed more light on finguring out the variables effecting carbon emission.

II. PROBLEM STATEMENTS

- 1) How can we develop high performing models to predict an individual's carbon emission based on their behavior?

- 2) How can we develop computationally efficient models to predict an individual's carbon emission based on their behavior?
- 3) How can we develop highly interpretable models to predict an individual's carbon emission based on their behavior?

III. SOLUTIONS

To address the problem of predicting an individual's carbon emissions, we implemented a comprehensive machine learning workflow designed to develop models that are not only high-performing, but also computationally efficient and interpretable. Our solution directly answers the questions posed in the problem statement through a multi-stage process:

1) *Enhanced Data Preprocessing and Feature Engineering*

We moved beyond naive data handling by implementing a sophisticated preprocessing pipeline. This involved structured ordinal encoding for categorical variables, consolidation of redundant features, and expanded one-hot encoding for multi-label fields. This structured approach was crucial for creating a rich, machine-readable feature set that significantly boosted model performance, as evidenced by the dramatic R^2 value increase in our Linear Regression and MLP models (Table III).

2) *Systematic Model Evaluation and Selection*

A rigorous benchmark between nine distinct machine learning models, ranging from simple linear regressors to complex deep learning and gradient-boosted tree architectures (Table I), were undergo. This comparative analysis allowed us to identify models that offered the most optimal trade-offs between performance and efficiency. The **CatBoost** model emerged as the superior choice, demonstrating state-of-the-art predictive capabilities.

3) Advanced Model Interpretation

To ensure our most accurate model was not a "black box," i.e., it produces a particular result with no demonstration of understanding of how it works, we employed SHapley Additive exPlanations (SHAP). This technique provides clear, actionable insights by quantifying the exact contribution of each behavioral feature to an individual's predicted carbon footprint (Fig. 5 & 6). This directly satisfies the requirement for a highly interpretable model.

Our final proposed solution is a **Catboost model** built upon an optimized feature engineering pipeline which yielded the highest performance while the model's inherent efficiency and the application of SHAP analysis ensure it meets all three objectives outlined in our problem statement.

IV. CONTRIBUTION

This paper makes several key contributions to the field of individual carbon footprint prediction, directly providing solutions to the core research questions:

1) A High-Performance Predictive Model

Our primary contribution is the development and validation of a CatBoost-based model that predicts an individual's carbon emissions with 99.2% accuracy (R^2 score). By systematically comparing nine different algorithms, we demonstrate that gradient-boosted trees, and CatBoost in particular, are exceptionally well-suited for this task, significantly outperforming traditional linear models, tree-based ensembles, and even standard neural networks as shown in Table III.

2) A Blueprint for Computationally Efficient Modelling

We address the need for efficiency by analyzing the performance of models known for their low computational overhead. Our results highlight that lightweight yet powerful models like LightGBM ($R^2 = 0.9765$) and **CatBoost** provide -near-cutting-edge performance without the intensive resource requirements of deep learning models like an MLP. This demonstrates that high accuracy and computational efficiency are not mutually exclusive, making our findings applicable to resource-constrained or real-time applications.

3) A Framework for High Model Interpretability

We successfully solve the challenge of interpretability in complex models by integrating SHAP Analysis. Our work moves beyond simple performance metrics to provide a clear explanation of why the model makes its predictions. The SHAP analysis (Fig. 5) definitively identifies transportation and consumption habits as the primary drivers of emissions, providing actionable insights for consumers, policymakers, and corporations. This contribution ensures our model is not only accurate, but also a valuable tool for driving meaningful behavioral change.

V. LITERATURE REVIEW

Machine learning (ML) models have been extensively applied to estimate CO_2 emissions at multiple scales—from regional and sectoral to the increasingly vital individual-consumer level. The recent systematic review by [7] synthesizes performance metrics across model families and serves as a key benchmark for this work. Similarly, [2] presents a broad survey of ML applications across CO_2 prediction contexts, offering insight into modeling techniques, loss function design, and model suitability for different problem types. These findings guide our modeling efforts to predict individual-level emissions using behavioral features.

A. Model Comparison and Behavior Fit

Linear and Kernel Models: Linear regression serves as a baseline model in many studies, including [7], where it achieved an average R^2 of 46%. While interpretable and computationally light, linear models struggle to model complex consumer behavior and exhibit high bias. Support Vector Regression (SVR), especially when combined with preprocessing methods like grey modeling or PCA, captures moderate nonlinearities and improves performance over linear baselines. [2] notes that GM-SVR hybrid models outperform standalone SVR in multi-provincial prediction tasks, with better generalization and reduced MAE.

k-Nearest Neighbors (k-NN): Non-parametric and easy to implement, k-NN improves predictive accuracy ($R^2 \approx 68\%$) and reduces MAPE to 37% as shown in [7], though it performs poorly with sparse, high-dimensional behavioral data due to the curse of dimensionality.

Tree-Based Ensembles (Random Forest, Boosting): Random Forest (RF) and Gradient Boosting models are particularly well-suited to heterogeneous and rule-driven features typical of individual behaviors. RF provides robust generalization, with [7] reporting R^2 values around 75% and MAPE near 33%. Boosted methods such as AdaBoost achieve the highest reported performance, with R^2 up to 92% and MAPE as low as 15%. [2] corroborates that ensemble methods remain state-of-the-art across varied emission domains, especially with categorical-heavy or mixed-type features.

Deep Models (ANN, LSTM, DPRNN): Feedforward neural networks and sequential models such as LSTM and DPRNN can capture complex and time-dependent behavioral patterns. While [2] demonstrates that LSTM outperforms traditional models on daily emission trends (achieving $R^2 \approx 0.98$), [7] cautions that deep learning methods may lack interpretability and require significantly more computational resources.

B. Model Behavior vs. Data Behavior

- **Linear/SVR:** Good for low-noise, low-complexity behaviors with strong linear or near-linear relationships.
- **k-NN:** Useful when similar individuals yield similar emissions, but less reliable under sparse conditions.
- **RF/AdaBoost:** Highly adaptive to modal, nonlinear, and mixed-type behavior; robust against overfitting.

- **LSTM/DPRNN:** Best for sequential data (e.g., temporal lifestyle tracking), but sensitive to overfitting and resource intensive.

C. Computational Cost and Interpretability

- **Linear Models:** Low cost, high interpretability.
- **SVR and k-NN:** Moderate cost; SVR interpretability medium, k-NN low.
- **RF/AdaBoost:** Moderate to high training cost; good interpretability via feature importance.
- **LSTM/DPRNN:** High training/inference cost; low interpretability without post hoc methods.

D. Loss Function and Evaluation Metrics

Most models employ Mean Squared Error (MSE) or Mean Absolute Percentage Error (MAPE). [7] emphasizes the use of Root Mean Squared Log Error (RMSLE) due to its effectiveness in handling skewed data and underprediction penalties. [2] also shows that hybrid loss functions, such as weighted MSE with structural penalties, can enhance deep model performance in daily and provincial predictions.

E. Conclusion and Benchmark Insight

In predicting individual behavior-based CO₂ emissions, ensemble models like RF and AdaBoost demonstrate high R^2 values (up to 92%) and low MAPE (as low as 6.3%) [7]. Deep learning models such as LSTM and DPRNN further improve prediction for time-varying behaviors [2], though with higher costs. Linear and kernel models offer efficient but less expressive baselines. As our work focuses on personalized and interpretable predictions, tree-based ensembles and LSTM variants form a promising modeling base.

VI. METHODOLOGY

A. Dataset Description

The Carbon Emission dataset comprises 10,000 individual records with 20 features, designed to support the development of predictive models for estimating CO₂ emissions. Seven of the features are numerical and the remaining thirteen are categorical. These features offer a comprehensive view of the lifestyle and demographic factors that influence an individual's carbon footprint.

There are no duplicate records in the dataset, and all values fall within their anticipated ranges. Additionally, the numerical variables and twelve out of thirteen categorical variables do not exhibit missing values. However, the categorical variable Vehicle Type is an exception, containing 6,721 missing values, which accounts for approximately 67 of the total records. This issue will be addressed during preprocessing, possibly by introducing a "None" category for respondents who do not own vehicles or by applying an appropriate imputation method, prior to exploratory analysis and model training.

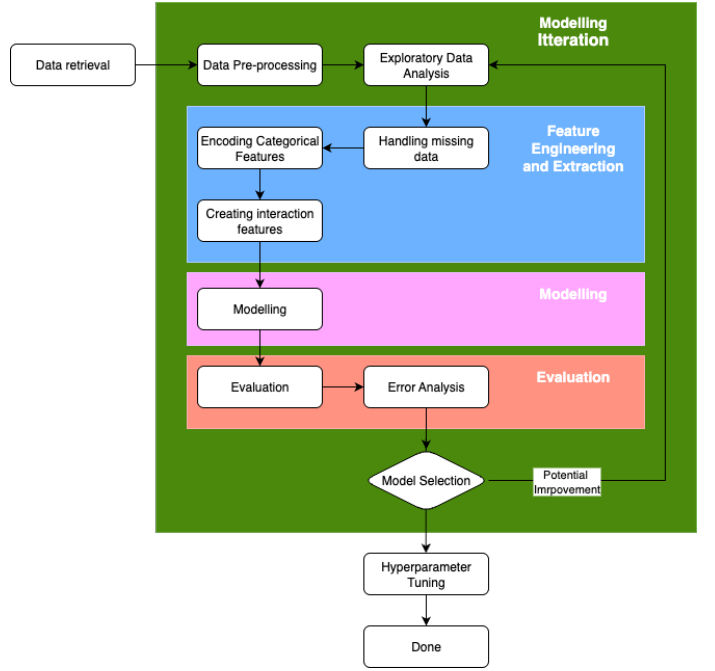


Fig. 1: Research Flowchart.

B. Method of Analysis

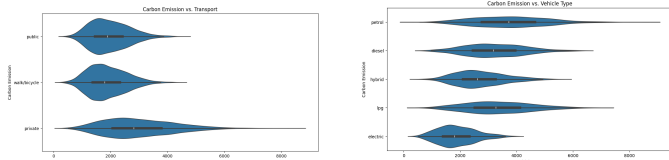
Our research were conducted in, mainly, 3 phases: (1) Pre-processing, (2) Modelling Iteration, and (3) Post-processing. The pre-processing phase includes retrieving and cleansing the data. Exploratory Data Analysis (EDA) were also underwent in this phase. Afterwards, we gather insights from our EDA and bring them to the modelling iteraion phase where we conducted feature engineering and extraction, modelling, and evaluation. Through this process, we developed features based on analysing a correlation heatmap which lead us to develop more effective features by dropping, encoding, and consolidating existing features. Finally, we iterate the same process until a superior model was decided which we then continue to improve that model by conducting hyperparameter tuning.

C. Exploratory Data Analysis

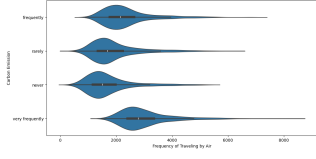
We explored the data's behavior and how we may exploit it for the best model performance and understanding. We plotted each feature's value distribution with respect to the Carbon Emmision typeface feature. We found that most features barely exhibit any significant pattern that may tell some relations to the Carbon Emmision typeface feature. Here, we present features having quite significant pattern.

VII. DATA ANALYSIS AND FEATURE ENGINEERING

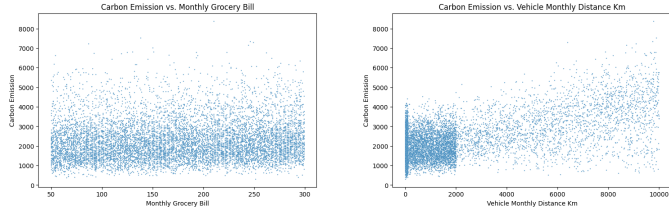
A correlation heatmap was constructed to evaluate the linear relationships among numerical and encoded categorical features with respect to carbon emissions. The heatmap (Fig. ??) revealed that most variables have weak pairwise correlations (generally $|r| < 0.2$), suggesting low linear dependence. As such, models relying on linear assumptions—like Linear



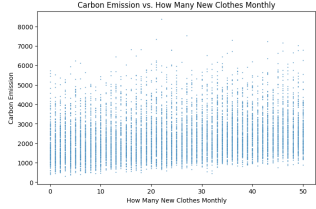
(a) Carbon Emission by Transport (b) Carbon Emission by Vehicle Type



(c) Carbon Emission by Frequency of Air Travel



(d) Carbon Emission vs. Monthly Grocery Bill (e) Carbon Emission vs. Vehicle Monthly Distance



(f) Carbon Emission by New Clothes Bought Monthly

Fig. 2: Carbon emissions distribution across selected individual behavioral factors.

Regression—may perform poorly, while tree-based and ensemble models, such as CatBoost and Gradient Boosting, are more suitable due to their ability to capture non-linear and interaction effects.

A. Feature Engineering

Based on the correlation insights and exploratory analysis, the following transformations were applied:

- **Dropped Features:** How Long TV PC Daily Hour, Energy Efficiency, Diet Type, Shower Frequency, Social Activity Frequency, Sex.
- **Ordinal Encoding:**
 - **Body Type:** Obese = 2, Overweight = 1, Underweight = 0
 - **Vehicle Type:** Petrol = 3, Hybrid = 2, Electric = 1, LPG = 0

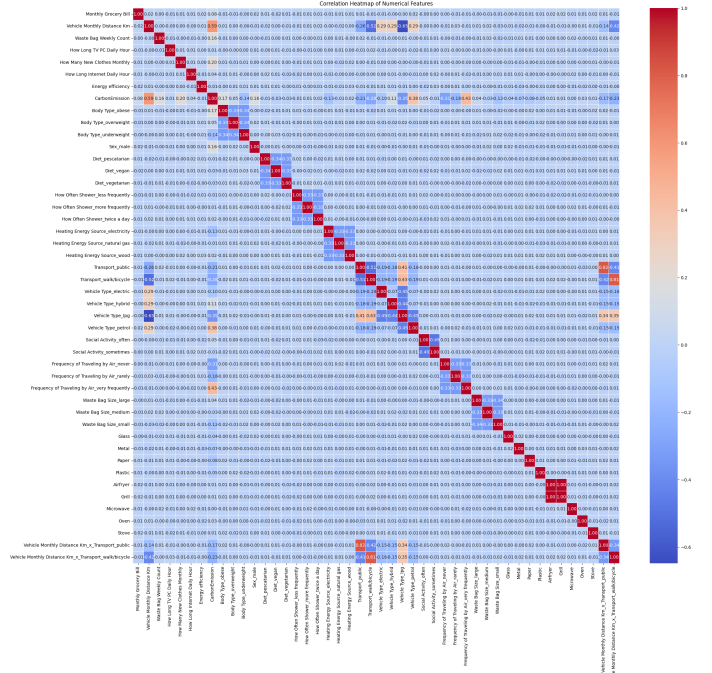


Fig. 3: Correlation Heatmap of Pre-processed Data

– **Other Features:** Frequency of Travelling by Air, Waste Bag Size, Transport Type, and Heating Energy Source were similarly ordinal-encoded.

- **Feature Consolidation:** Similar appliance features (e.g., Air Fryer and Grill) were merged to reduce redundancy and enhance generalizability.

B. Model Evaluation

We trained and evaluated several machine learning models using three standard metrics: Mean Absolute Error (MAE), Mean Squared Error (MSE), and coefficient of determination (R^2). The dataset was processed through consistent preprocessing pipelines. Table I summarizes the results:

TABLE I: Model Performance Comparison

Model	MAE	MSE	R^2
Linear Regression	351.55	250669.67	0.7589
Support Vector Regressor	666.52	760429.90	0.2686
Decision Tree	371.91	240803.89	0.7684
Random Forest	255.52	113635.95	0.8907
Gradient Boosting	226.02	90392.06	0.9131
XGBoost	223.64	92815.51	0.9107
LightGBM	211.13	78656.80	0.9243
CatBoost	199.49	70312.31	0.9324
Neural Network (MLP)	350.26	253447.06	0.7562

Algorithm 1 Carbon Emission Modeling Workflow from Heatmap Analysis

```
1: Input: Raw behavioral dataset with individual features
2: Output: Trained carbon emission prediction model
3: // Step 1: Preprocessing
4: Drop features: TV/PC Usage, Energy Efficiency,
   Diet Type, Shower Frequency, Social
   Activity Frequency, Sex
5: Encode Body Type: Obese → 2, Overweight → 1, Underweight → 0
6: Encode Vehicle Type: Petrol → 3, Hybrid → 2, Electric → 1, LPG → 0
7: Encode categorical features: Air Travel Frequency, Waste Bag Size, Transport Type, Heating Source
8: Merge overlapping appliance features: Air Fryer, Grill
9: // Step 2: Model Training
10: for all models ∈ {Linear Regression, SVR, Decision Tree, Random Forest, Gradient Boosting, XGBoost, LightGBM, CatBoost, MLP} do
11:   Train model on 80% of data
12:   Validate model on 20% of data
13:   Evaluate metrics: MAE, MSE,  $R^2$ 
14: end for
15: // Step 3: Model Selection
16: Select model with highest  $R^2$  and lowest MAE/MSE
17: Return: Best model (CatBoost)
```

Algorithm 2 Enhanced Data Preprocessing and Feature Encoding

```
1: Input: Raw dataset df_test
2: Output: Transformed and encoded dataset
3: // Step 1: Drop Irrelevant Features
4: Drop columns: How Long TV PC Daily Hour, Energy efficiency, Diet, How Often Shower, Social Activity, Sex
5: // Step 2: Encode Categorical Features with Ordinal Mappings
6: Encode Body Type: underweight → 0, normal → 1, overweight → 2, obese → 3
7: Encode Vehicle Type: lpg → 0, electric → 1, hybrid → 2, petrol/diesel → 3
8: Fill missing values in Vehicle Type with -1
9: Encode Frequency of Traveling by Air: never → 0, rarely → 1, frequently → 2, very frequently → 3
10: Encode Waste Bag Size: small → 0, medium → 1, large → 2, extra large → 3
11: Encode Heating Energy Source: wood/natural gas → 0, electricity → 1, coal → 2
12: Encode Transport: walk/bicycle → 0, public → 1, private → 2
13: // Step 3: Parse Multi-Label Columns into Lists
14: Convert string lists in Recycling and Cooking_With into Python list objects
15: // Step 4: One-Hot Encode Multi-Label Features
16: Generate dummy variables from Recycling and concatenate with dataset
17: Drop original Recycling column
18: Generate dummy variables from Cooking_With and concatenate with dataset
19: Drop original Cooking_With column
20: // Step 5: Feature Consolidation
21: if Airfryer and Grill exist in dataset then
22:   Create binary column Airfryer_and_Grill using logical OR
23:   Drop columns Airfryer and Grill
24: end if
25: Return: Cleaned, encoded dataset df_test
```

C. Interpretation

CatBoost outperformed all other models, achieving the lowest MAE and MSE and the highest R^2 score (0.9324). LightGBM and Gradient Boosting followed closely, showing strong capability in handling the dataset's mixed feature types and non-linear structure. In contrast, linear models and Support Vector Regression exhibited substantially weaker performance, reinforcing the conclusion drawn from the correlation heatmap analysis.

These results demonstrate the effectiveness of gradient-boosted tree-based models in capturing complex behavioral patterns associated with individual carbon emissions.

We also performed a more naive feature engineering, in the hopes that the models may pickup patterns by itself without extensive data wrangling

VIII. RESULTS AND DISCUSSIONS

A. Comparative Analysis of Model Performance

To assess the effectiveness of the enhanced preprocessing pipeline, we compared model performance before and after applying structured encodings and feature refinements. The updated preprocessing involved ordinal encoding for categorical variables (e.g., Body Type, Vehicle Type, Transport), consolidated binary indicators (e.g., Airfryer & Grill), and expanded one-hot encoding of multi-label fields (Recycling, Cooking_With).

Table II summarizes the R^2 score improvements across different machine learning models. The enhancements led to significant performance gains, especially for models sensitive to input representation quality.

TABLE II: Performance Comparison Before vs. After Preprocessing

Model	Old R^2	New R^2	ΔR^2
Linear Regression	0.7589	0.9370	+0.1781
Gradient Boosting	0.9131	0.9502	+0.0371
Support Vector Regressor	0.2686	0.2665	−0.0021
Neural Network (MLP)	0.7562	0.9732	+0.2170
Decision Tree	0.7684	0.7973	+0.0289
Random Forest	0.8907	0.9157	+0.0250
XGBoost	0.9107	0.9729	+0.0622
LightGBM	0.9243	0.9765	+0.0522
CatBoost	0.9324	0.9920	+0.0596

Observations

- **CatBoost** consistently outperformed all other models. After preprocessing, its R^2 increased from 0.9324 to 0.9920, while its MAE dropped from 199.49 to 66.37, confirming its strength in handling structured categorical inputs.
- **Neural Network (MLP)** exhibited the largest gain in R^2 (+0.217), showing that neural models are particularly sensitive to input noise and benefit greatly from refined representations.
- **Linear Regression**, despite its simplicity, improved remarkably by nearly 18% in R^2 —highlighting the power of proper encoding even for linear methods.
- **Support Vector Regressor (SVR)** showed negligible improvement and remained the poorest performer, suggesting incompatibility with the feature space or scale.
- **Tree-based models** like LightGBM and XGBoost also improved moderately, with final R^2 scores exceeding 0.97, close behind CatBoost.

B. Model Interpretability via SHAP Analysis

To interpret the predictions of the final CatBoost model ($R^2 = 0.9920$), we employed SHapley Additive exPlanations (SHAP), which attribute the contribution of each feature to individual predictions. Figure 4 presents the global feature importance, while Figure 5 shows the distribution of SHAP values for each feature.

1) Global Feature Importance

The bar plot in Figure 4 displays the average magnitude of SHAP values, indicating each feature’s contribution to model output, regardless of direction.

- **Vehicle Monthly Distance Km** emerged as the most influential feature by a large margin, suggesting a strong positive relationship between longer vehicle usage and higher carbon emissions.
- **Air Travel Frequency** (particularly very frequently and never) ranked highly, reflecting the emission impact of long-distance transport behavior.
- **How Many New Clothes Monthly** and **Sex (male)** also showed substantial influence, possibly capturing consumption and lifestyle-driven emissions.
- **Waste Bag Weekly Count** and **Body Type (obese)** were moderate contributors, potentially linking dietary and household waste behavior to emission levels.
- Less influential but non-negligible features included Vehicle Type, Waste Bag Size, Social Activity Frequency, and Heating Energy Source.

2) Feature-Wise Impact Distribution

The beeswarm plot in Figure 5 illustrates how individual feature values affect the model output.

- For **Vehicle Monthly Distance Km**, high values (shown in red) are strongly associated with positive SHAP values—indicating greater predicted emissions.
- **Air Travel Frequency (very frequently)** similarly shows a positive influence when high, while the `never` category contributes to negative SHAP values (lower predicted emissions).
- Interestingly, **Electric Vehicle Type** has a distribution centered around zero, suggesting a neutral or weak mitigating effect, while **Petrol** vehicles slightly increase emissions.
- **Sex (male)** and **Clothing Consumption** show diverse effects, indicating interaction with other behavioral factors.

3) Outcomes

The SHAP analysis confirms that carbon emissions are most heavily influenced by transportation behaviors—both in frequency and mode—as well as air travel and consumer lifestyle indicators. This interpretable output supports the model’s predictions and can inform personalized feedback strategies in sustainability applications.

Conclusions

These results demonstrate that data preprocessing has a substantial impact on predictive performance, particularly for non-linear and ensemble-based models. Gradient-boosted frameworks—especially CatBoost—excel in capturing behavioral-carbon relationships when categorical and multi-label data are well-engineered.

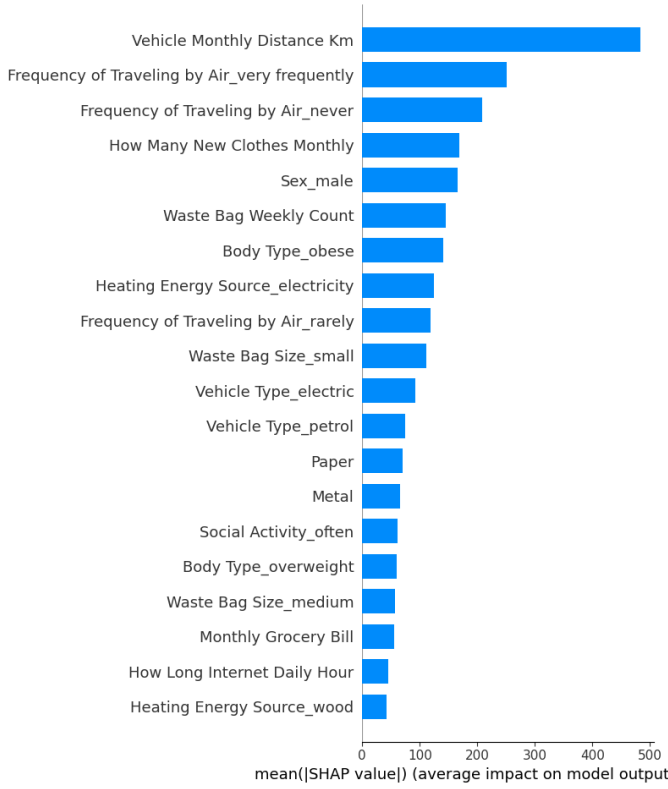


Fig. 4: Mean SHAP values for top features in the CatBoost model.



Fig. 5: SHAP value distribution by feature (beeswarm plot).

C. Discussions

1) Modelling

Select and train various regression models (such as Linear Regression, Gradient Boosting, SVM, Neural Networks, Tree-based models, XGBoost, LightGBM, CatBoost, etc.) to predict ‘CarbonEmission’ and compare their performance.

2) Results and Interpretations

Another modeling approach was centered around a Stacking Regressor. This ensemble technique is a widely-used method chosen to combine the predictive power of multiple diverse, high-performing algorithms to achieve better performance than any single model could on its own [6]. The base learners included Random Forest, XGBoost, and CatBoost, which are considered state-of-the-art models known for their accuracy and robustness. By using a Ridge Regressor as the final meta-learner, the model aggregates the predictions from the base learners to produce a final, more refined estimate. This hierarchical structure is designed to capture complex patterns in the data that a single model might miss, thereby enhancing overall predictive performance.

The model’s performance was evaluated using 5-fold cross-validation and a hold-out dataset, with the Coefficient of Determination (R^2) and Root Mean Squared Error (RMSE) as key metrics. The results were outstanding, showing a Coefficient of Determination (R^2) of 0.991 for both the cross-validation and the hold-out sets. This indicates the model explains 99.1% of the variance in carbon emissions and, more importantly, generalizes exceptionally well to new data without overfitting. The consistency between the training ($R^2 = 0.994$) and validation scores further confirms its robustness. Additionally, the low hold-out RMSE of 97.694 signifies a very small prediction error, reinforcing the model’s high accuracy and stability.

Overall, we achieved the best model performance from this research by applying the naive approach of feature engineering with performance as follows:

TABLE III: Model Performance Metrics (After Enhanced Preprocessing)

Model	MAE	MSE	R^2
Linear Regression	168.04	65542.20	0.9370
Gradient Boosting	169.69	51754.10	0.9502
Support Vector Regressor	667.32	762675.95	0.2665
Neural Network (MLP)	122.80	27883.48	0.9732
Decision Tree	343.64	210728.16	0.7973
Random Forest	229.06	87627.97	0.9157
XGBoost	127.54	28190.97	0.9729
LightGBM	117.75	24441.53	0.9765
CatBoost	66.37	8297.93	0.9920

REFERENCES

- [1] Alnuaimi, H. S., Acquaye, A., and Mayyas, A. (2025). Machine learning applications for carbon emission estima-

- tion. *Resources Conservation Recycling Advances*, page 200263.
- [2] Hong, Y., He, P., Liu, W., Zhang, Z., Zhou, X., Tang, J., Han, Z., and Song, J. (2025). A systematic review of carbon emission prediction using machine learning: Model comparison, interpretability, and hybrid optimization. *Cleaner Energy Systems*, 6:100104.
 - [3] Matasci, C., Gauch, M., Böni, H., and Wäger, P. (2021). The influence of consumer behavior on climate change: the case of switzerland. *Sustainability*, 13(5):2966.
 - [4] Nations”, U. (n.d.). The 17 goals — sustainable development.
 - [5] Ritchie, H. and Roser, M. (2020). *CO₂ emission. Our World in Data*.
 - [6] Sagi, O. and Rokach, L. (2018). Ensemble learning: A survey. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 8(4):e1249.
 - [7] Serafeim, G. and Caicedo, G. V. (2022). Machine learning models for prediction of scope 3 carbon emissions. Technical Report 22-080, Harvard Business School. Working Paper.
 - [8] Thøgersen, J. (2021). Consumer behavior and climate change: consumers need considerable assistance. *Current Opinion in Behavioral Sciences*, 42:9–14.