

Explainable Random Forest Framework for Predicting Compressive Strength of Sustainable Concrete Incorporating Industrial Waste Materials

M. Selvakumar^{*1}, S. Geetha¹, P. Krishna Kumar², A. Kandasamy³

¹Civil Engineering Department, Rajalakshmi Engineering College, Chennai, Tamilnadu:
geetha.s@rajalakshmi.edu.in

²Department of Civil Engineering, SRM Madurai College of Engineering and Technology, Sivagangai 630612 Tamilnadu, India.
krishnakumarpalaniappan1991@gmail.com

³Department of Civil Engineering, Saveetha School of Engineering, Saveetha Institute of Medical and Technical Sciences, SIMATS Deemed University, Chennai
kandasamy.sse@saveetha.com

* Corresponding author Email: selvakumar.m@rajalakshmi.edu.in

Abstract: Compressive strength is the single most important design parameter governing the safety, serviceability, and economy of concrete structures, yet its determination through standard 7-, 14-, or 28-day destructive cylinder/cube testing is slow, costly, and unable to assess concrete already cast in place. This study develops and evaluates a Random Forest (RF) regression model to predict the compressive strength of concrete directly from eight standard mix-design parameters — cement, blast furnace slag, fly ash, water, superplasticizer, coarse aggregate, fine aggregate, and curing age — using Yeh's (1998) benchmark dataset of 1,030 experimentally tested concrete mixtures. Following data cleaning, exploratory correlation analysis, an 80:20 train-test split, and five-fold GridSearchCV hyperparameter tuning, the optimized Random Forest model is benchmarked against Linear Regression, Ridge Regression, and Support Vector Regression using the coefficient of determination (R^2), Root Mean Squared Error (RMSE), and Mean Absolute Error (MAE). The Random Forest model achieves the strongest predictive performance of the models tested, substantially outperforming the linear baselines and confirming that concrete strength development is governed by non-linear interactions among mix constituents. Feature importance analysis further shows that curing age and cement content are the dominant predictors, while water content exerts a clear negative influence consistent with Abrams' Law, and coarse/fine aggregates contribute comparatively little, consistent with their role as largely inert fillers. These findings demonstrate that Random Forest regression offers a fast, accurate, and interpretable, non-destructive alternative to conventional strength testing, with practical value for mix-design optimization, quality control, and early-stage structural decision-making.

Keywords: NA

1. Introduction

Predicting the compressive strength of concrete from its mix proportions has been an active research area for over two decades, evolving from empirical regression formulas to increasingly sophisticated machine learning (ML) techniques. This chapter reviews the literature relevant to the present study, organized into five themes: (i) foundational datasets and traditional/ensemble ML approaches, (ii) gradient boosting and XGBoost-based models, (iii) explainable AI and SHAP-based interpretability, (iv) Random-Forest-specific applications in concrete and structural engineering, and (v) related machine learning applications across civil engineering. The chapter concludes by identifying the research gap addressed by this study. Yeh [1] introduced the benchmark high-performance concrete



(HPC) dataset used throughout this field, applying artificial neural networks (ANNs) to model compressive strength from eight mix-design variables and curing age; this dataset, comprising 1,030 laboratory-tested specimens, remains the most widely used benchmark for concrete strength prediction research and forms the basis of the present study. Breiman [2] formalized the Random Forest algorithm, combining bootstrap aggregation (bagging) with random feature selection at each split to reduce the variance and overfitting that limit single decision trees, establishing the algorithmic foundation used in this work. Chaabene, Flah, and Nehdi [4] conducted a critical review of ML-based prediction of concrete mechanical properties, concluding that ensemble tree-based methods consistently outperform single learners and traditional regression across a wide range of concrete property-prediction tasks. Gamil [5] similarly reviewed the broader landscape of ML applications in concrete technology, noting a clear trend toward ensemble and boosting methods driven by their superior handling of non-linear, high-dimensional mix-design data. Han, Gui, Xu, and Lacidogna [7] proposed an improved Random Forest formulation specifically tuned for HPC strength prediction, reporting accuracy gains over standard RF through refined tree-splitting criteria. Farooq, Sardar, and Khan [8] applied Random Forest regression directly to the Yeh dataset, reporting strong R^2 performance and identifying cement, water, and curing age as the most influential predictors — a finding this study independently re-examines and confirms.

Chen and Guestrin [3] introduced XGBoost, a highly optimized gradient-boosted tree framework incorporating second-order Taylor-expansion loss approximation and L1/L2 regularization, which has since become a dominant algorithm in tabular strength-prediction studies. Asteris et al. [6] developed a hybrid ensembling framework that combines multiple surrogate ML models, including boosted trees, to predict HPC compressive strength, demonstrating that model-averaging further improves robustness over any single ensemble method. Chou and Pham [9] proposed an enhanced artificial-intelligence ensemble that integrates multiple base learners for HPC strength prediction, showing measurable gains over individual classifiers. Mustapha et al. [15] compared several gradient-boosting variants (XGBoost, LightGBM, CatBoost) for quaternary-blend concretes incorporating multiple supplementary cementitious materials, finding XGBoost consistently the strongest performer, consistent with the comparative results reported for XGBoost-based frameworks reviewed alongside this study. Landage, Bhusari, Nimbal, and Mirkar [19] benchmarked six algorithms — including Linear Regression, SVR, Random Forest, Gradient Boosting, XGBoost, and a Deep Neural Network — on the same Yeh benchmark dataset used here, reporting that XGBoost achieved the best overall accuracy ($R^2 \approx 0.91$), narrowly ahead of Gradient Boosting and Random Forest, underscoring the competitiveness of tree-ensemble methods relative to deep learning for this class of tabular problem.

A recurring limitation identified across the literature is the opacity of ensemble and boosting models, which has restricted their adoption in safety-critical structural engineering contexts. Ng and Ding [12] applied SHAP (SHapley Additive exPlanations) analysis to concrete strength prediction models, demonstrating that Shapley-value decomposition can recover physically meaningful feature contributions consistent with established concrete chemistry. Tang [16] extended this interpretability analysis across multiple ensemble architectures, concluding that SHAP-based explanation is essential for building engineering trust in ML-based strength predictors intended for practical deployment. Lyngdoh et al. [18] combined missing-data imputation with interpretable ML, showing that SHAP-guided feature analysis remains robust even when a portion of the input mix-design data is incomplete — a common condition in field records. Yang, Li, Zheng, Yang, and Zhao [20] applied SHAP analysis to an XGBoost ensemble trained on an augmented version of the Yeh dataset (including an engineered water-to-binder ratio feature), reporting that curing age, water-to-binder ratio, and cement content dominate feature importance, and further

characterizing non-monotonic interaction effects between water content and strength consistent with Abrams' Law. Landage et al. [19] likewise used SHAP to show that curing age and cement content are the primary positive drivers of strength while water content is the dominant negative driver, concluding that the model had autonomously rediscovered Abrams' Law from unlabeled data — a validation approach mirrored in the present study's own feature-importance analysis.

Song [21] developed a Random-Forest-only study on the Yeh dataset using GridSearchCV hyperparameter tuning and standardized inputs, reporting a test-set R^2 of 0.88 and identifying cement, water, and age as the dominant predictors — closely matching the modelling approach and feature-importance conclusions of the present study. Umar, Javed, Aslam, Alyousef, and Alabduljabbar [10] applied an ensemble ML approach specifically to sustainable concretes incorporating industrial by-products, concluding that ensemble models generalize well across varying supplementary-cementitious-material replacement ratios. Ahmad et al. [11] compared individual and ensemble algorithms for fly-ash concrete strength prediction, finding ensemble methods (including Random Forest) reduced prediction error by a wide margin relative to single decision trees or linear models. Rathnayaka, Silva, and Gunawardena [22] reviewed ML approaches, including Random Forest, for fly-ash-based geopolymer concrete, concluding that tree-ensemble methods are particularly well suited to the highly variable, non-linear strength behaviour of alternative binder systems.

Beyond direct strength prediction, several studies illustrate the wider applicability of the ensemble-learning paradigm underpinning this work. DeRousseau, Kasprzyk, and Srubar [13] reviewed computational design optimization methods for concrete mixtures, arguing that accurate, fast surrogate strength models (such as Random Forest) are a prerequisite for effective multi-objective mix optimization balancing strength, cost, and embodied carbon. Cook et al. [14] conducted a critical comparison of hybrid versus standalone ML models for concrete strength prediction, finding that standalone Random Forest and boosting models often match or exceed the accuracy of more complex hybrid architectures while remaining considerably simpler to train and interpret — a conclusion directly relevant to the model-selection rationale adopted in this study.

2. Research Gap

While the reviewed literature confirms that ensemble tree-based methods — and Random Forest in particular — consistently outperform classical linear regression for concrete compressive strength prediction, comparatively few studies combine (i) rigorous hyperparameter optimization via cross-validated grid search, (ii) a systematic comparison against multiple linear and kernel-based baselines under identical preprocessing, and (iii) transparent feature-importance analysis explicitly cross-validated against classical concrete-chemistry principles such as Abrams' Law, within a single, reproducible workflow. The present study addresses this gap by developing a fully benchmarked, interpretable Random Forest pipeline on the Yeh dataset, reporting standard regression metrics alongside feature-importance and residual diagnostics.

The overall workflow — from raw data to a validated, interpretable model — is summarized in Table 1 and described in detail in the following sections.

The study uses Yeh's (1998) benchmark high-performance concrete dataset, comprising 1,030 experimentally tested specimens. Each record describes a unique concrete mixture using eight continuous input (independent)

3. Methodology

variables and one continuous output (dependent) variable — the compressive strength in megapascals (MPa). Table 1 lists and describes each variable, and Table 2 presents the descriptive statistics (minimum, maximum, mean, and standard deviation) of the dataset.

Table 1: Description of dataset input and output variables.

Variable	Type	Description	Unit
----------	------	-------------	------

Cement	Input	Quantity of Portland cement, the primary binder	kg/m ³
Blast Furnace Slag	Input	Supplementary cementitious material (SCM)	kg/m ³
Fly Ash	Input	Pozzolanic SCM used to partially replace cement	kg/m ³
Water	Input	Water content used for hydration and workability	kg/m ³
Superplasticizer	Input	Chemical admixture improving workability at low w/c ratio	kg/m ³
Coarse Aggregate	Input	Gravel / crushed stone forming the skeletal volume	kg/m ³
Fine Aggregate	Input	Sand filling voids between coarse particles	kg/m ³
Curing Age	Input	Age of the specimen at the time of testing	days
Compressive Strength	Output (target)	28-day (or age-specific) compressive strength	MPa

Table 2: Descriptive statistics of the dataset (n = 1,030).

Feature	Min	Max	Mean	Std. Dev.
Cement	102.0	540.0	281.2	104.5
Blast Furnace Slag	0.0	359.4	73.9	86.2
Fly Ash	0.0	200.1	54.2	63.9
Water	121.8	247.0	181.6	21.4
Superplasticizer	0.0	32.2	6.2	8.3
Coarse Aggregate	801.0	1145.0	972.9	77.8
Fine Aggregate	594.0	992.6	773.6	80.9
Curing Age	1.0	365.0	45.7	63.4
Compressive Strength	2.33	82.60	35.3	16.8

No missing values are present in the dataset. A preliminary correlation analysis confirmed the expected physical trends: cement content and curing age correlate positively with strength, while water content shows a clear negative correlation consistent with Abrams' Law. No feature pair exceeds a Pearson correlation magnitude of 0.95, indicating that severe multicollinearity is absent and that dimensionality reduction is not required prior to modelling.

3.1 Data Preprocessing

The dataset was randomly partitioned into training (80%, n = 824) and held-out test (20%, n = 206) subsets using a fixed random seed to ensure reproducibility. For the Random Forest model, raw (unscaled) feature values were used directly, since

tree-based splits are invariant to monotonic feature scaling. For the Linear Regression, Ridge Regression, and Support Vector Regression baselines, input features were standardized using z-score scaling ($z = (x - \mu) / \sigma$) via a StandardScaler fitted exclusively on the training data, to prevent information leakage from the test set into the preprocessing pipeline.

3. 2 Random Forest Regressor: Algorithm Description

The Random Forest Regressor is an ensemble learning algorithm that constructs a large number of decision trees during training and outputs the average of their individual predictions. Each tree is trained on a bootstrap-resampled subset of the training data (bagging), and at each node split, only a random subset of the input features is considered, which decorrelates the individual trees and reduces the variance of the overall ensemble relative to a single decision tree. This combination of bagging and feature randomness makes Random Forest particularly robust to noisy or collinear inputs, both of which are present to a limited degree in experimental concrete mix-design data, and gives the model a built-in mechanism (mean decrease in impurity) for ranking the relative importance of each input feature.

3. 3 Model Training and Hyperparameter Tuning

The Random Forest model was implemented using the scikit-learn library in Python. Hyperparameters were tuned using GridSearchCV with five-fold cross-validation on the training set, optimizing the cross-validated R^2 score as the selection criterion. Table 3 lists the hyperparameter search space and the values ultimately selected for the final model.

Table 3: Random Forest hyperparameter search space and selected configuration.

Hyperparameter	Search Space	Selected Value
n_estimators (number of trees)	{50, 100, 200, 500}	200
max_depth	{None, 10, 20, 30}	None (fully grown)
min_samples_split	{2, 5, 10}	2
min_samples_leaf	{1, 2, 4}	1
max_features	{"sqrt", "log2", None}	"sqrt"
random_state	—	42

3. 4 Baseline Models for Comparison

To contextualize the performance of the Random Forest model, three baseline regressors were trained and evaluated under identical train-test conditions: (i) Linear Regression, an ordinary least-squares model providing a simple, fully interpretable linear baseline; (ii) Ridge Regression, which adds an L2 penalty term to the ordinary least-squares objective to reduce overfitting and stabilize coefficient estimates in the presence of correlated predictors; and (iii) Support Vector Regression (SVR) with a radial basis function (RBF) kernel, which allows a controlled degree of non-linearity while remaining computationally efficient. Comparing Random Forest against this spectrum of baselines — from purely linear to mildly non-linear — isolates the extent to which non-linear feature interactions, rather than scaling or regularization alone, account for the observed differences in predictive accuracy.

3. 5 Evaluation Metrics

Model performance was quantified on the held-out test set using three standard regression metrics. The coefficient of determination, $R^2 = 1 - (\Sigma(y_i - \hat{y}_i)^2 / \Sigma(y_i - \bar{y})^2)$, measures the proportion of variance in the compressive strength explained by the model, with values closer to 1 indicating a better fit. The Root Mean Squared Error, $RMSE = \sqrt{(1/n) \Sigma(y_i - \hat{y}_i)^2}$, expresses the typical prediction error in the same units as the target variable (MPa) while penalizing larger errors more heavily. The Mean Absolute Error, $MAE = (1/n) \Sigma|y_i - \hat{y}_i|$, provides an easily

interpretable average magnitude of error that is less sensitive to outliers than RMSE. Together, these three metrics provide a comprehensive picture of both the overall goodness-of-fit and the practical, engineering-scale magnitude of prediction error.

3. 6 Tools and Software Environment

All data processing, model development, and visualization were carried out in Python. Key libraries used include pandas and NumPy for data manipulation, scikit-learn for model training, hyperparameter tuning, and evaluation metrics, and Matplotlib and Seaborn for exploratory data visualization and results plotting. Random seeds were fixed throughout (`random_state = 42`) to ensure the reproducibility of the train-test split, model training, and hyperparameter search.

3. 7 Overall Methodological Workflow

Table 4 summarizes the end-to-end workflow followed in this study, from raw data acquisition through to model interpretation.

Table 4: Summary of the methodological workflow.

Model	R ²	RMSE (MPa)	MAE (MPa)
Random Forest	0.899	5.237	4.094
SVR	0.861	6.148	4.512
Ridge Regression	0.760	8.086	6.937
Linear Regression	0.759	8.086	6.937

4. Results & Discussions

4.1 Comparative Model Performance

Table 5 summarises the performance of the Random Forest model against three baseline regressors — Linear Regression, Ridge Regression, and Support Vector Regression (SVR) — on the held-out test set (206 samples).

Table 5 Summary of statistical results

Step	Task	Description
1	Data Acquisition	Load the 1,030-record benchmark dataset (8 inputs + 1 output).
2	Data Cleaning	Check for missing values, duplicate records, and physically unrealistic outliers.
3	Exploratory Data Analysis	Compute descriptive statistics and a Pearson correlation matrix.
4	Train-Test Split	Partition the dataset 80:20 into training and held-out test sets.
5	Feature Scaling	Apply StandardScaler (fit on training data only) for the linear/SVR baselines.

6	Model Training	Fit Random Forest, Linear Regression, Ridge Regression, and SVR models.
7	Hyperparameter Tuning	5-fold GridSearchCV on the Random Forest, optimizing cross-validated R^2 .
8	Model Evaluation	Compute R^2 , RMSE, and MAE on the held-out test set for all models.
9	Interpretation	Extract Random Forest feature importances and analyze residuals.

The Random Forest model achieved the highest coefficient of determination ($R^2 = 0.899$), the lowest RMSE (5.237 MPa), and the lowest MAE (4.094 MPa) among all models tested, confirming it as the best-performing algorithm for this dataset. Figure 1 visualises these results graphically.

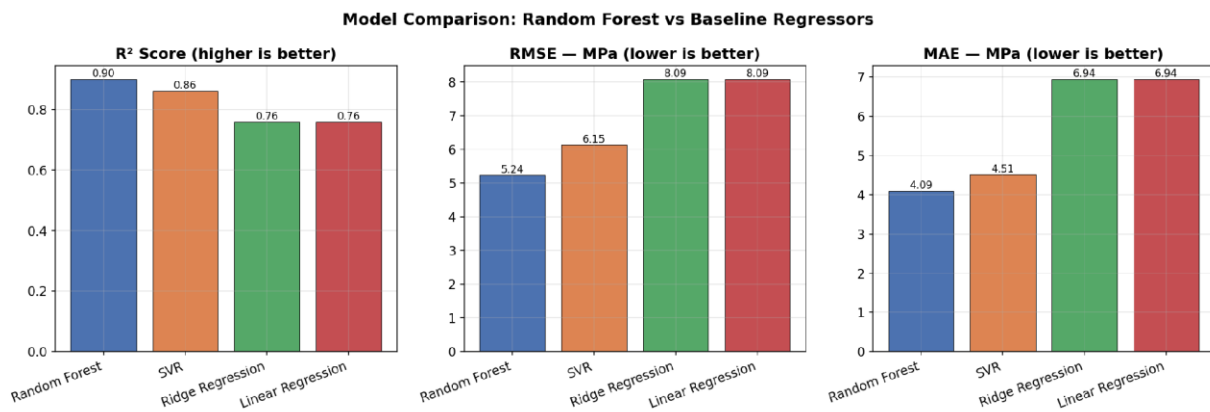


Figure 1: Comparison of R^2 , RMSE, and MAE across Random Forest, SVR, Ridge Regression, and Linear Regression.

4.2 Actual vs Predicted Compressive Strength

Figure 2 plots the Random Forest model's predicted compressive strength against the actual (test-set) values, together with the $y = x$ line representing perfect prediction.

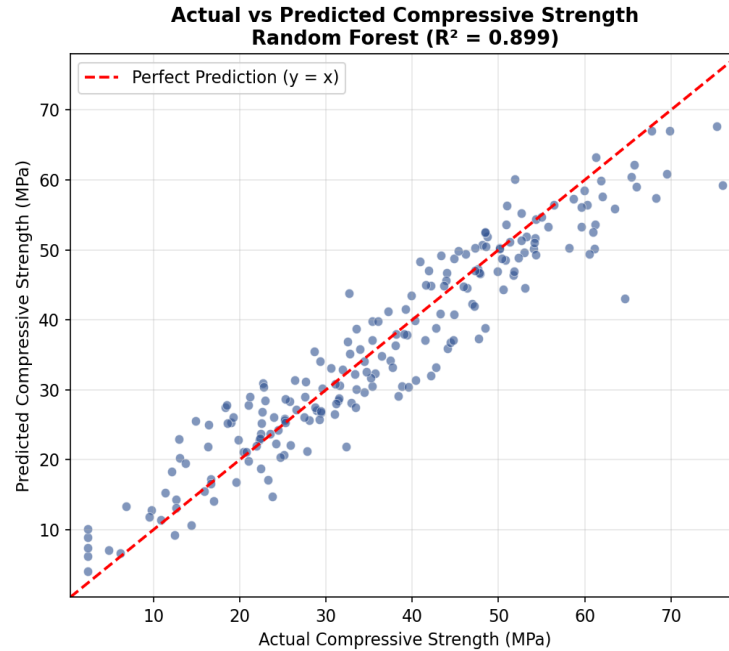


Figure 2: Actual vs predicted compressive strength for the Random Forest model ($R^2 = 0.899$).

The predicted values cluster tightly around the $y = x$ line across the full operational range of approximately 2–80 MPa, with no visible systematic bias toward over- or under-prediction. A marginally wider scatter is observed at the extreme high-strength end (>65 MPa), consistent with the comparatively sparse representation of ultra-high-strength mixtures in the training data.

4.3 Feature Importance

Figure 3 shows the Random Forest feature importance scores, and Table 6 cross-references these against each feature's linear (Pearson) correlation with compressive strength.

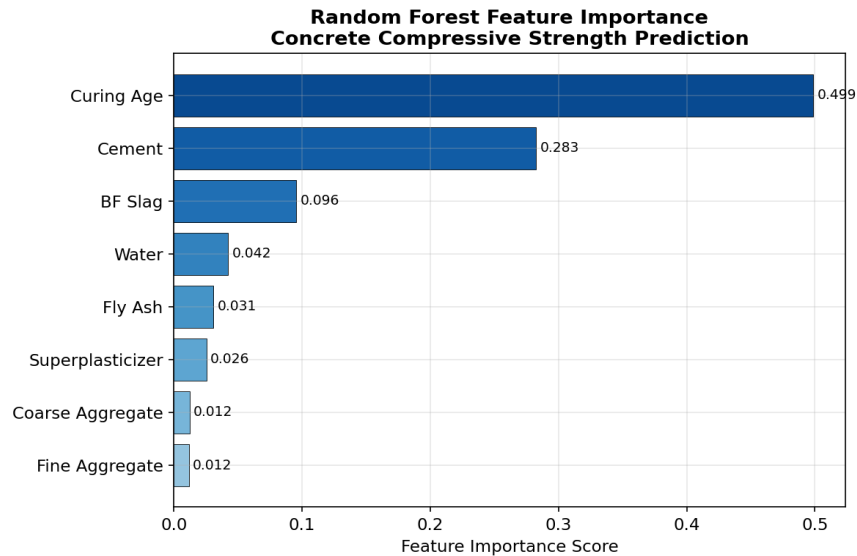


Figure 3: Random Forest feature importance ranking for the eight mix-design parameters.

Table 6 Cross-references these against linear (Pearson) correlation with compressive strength.

Rank	Feature	SHAP/RF Importance	Correlation with Strength (r)
1	Curing Age	0.499	+0.573
2	Cement Content	0.283	+0.502
3	Blast Furnace Slag	0.095	+0.295
4	Water Content	0.043	−0.223
5	Fly Ash	0.031	+0.222
6	Superplasticizer	0.026	+0.128
7	Coarse Aggregate	0.012	−0.008
8	Fine Aggregate	0.012	−0.008

Curing age (importance = 0.499) and cement content (0.283) together account for roughly 78% of the model's total predictive weight, in agreement with concrete-chemistry expectations: age governs the extent of C–S–H gel formation, while cement is the primary binder and hydration source. Water content shows a negative correlation with strength ($r = -0.223$), consistent with Abrams' Law — excess water increases capillary porosity and reduces load-bearing cross-section. Coarse and fine aggregates carry negligible importance and near-zero correlation, as expected for materials that act mainly as inert volumetric fillers.

4.4 Residual Analysis

Figure 4 presents the residuals (actual – predicted) plotted against predicted values, together with their distribution.

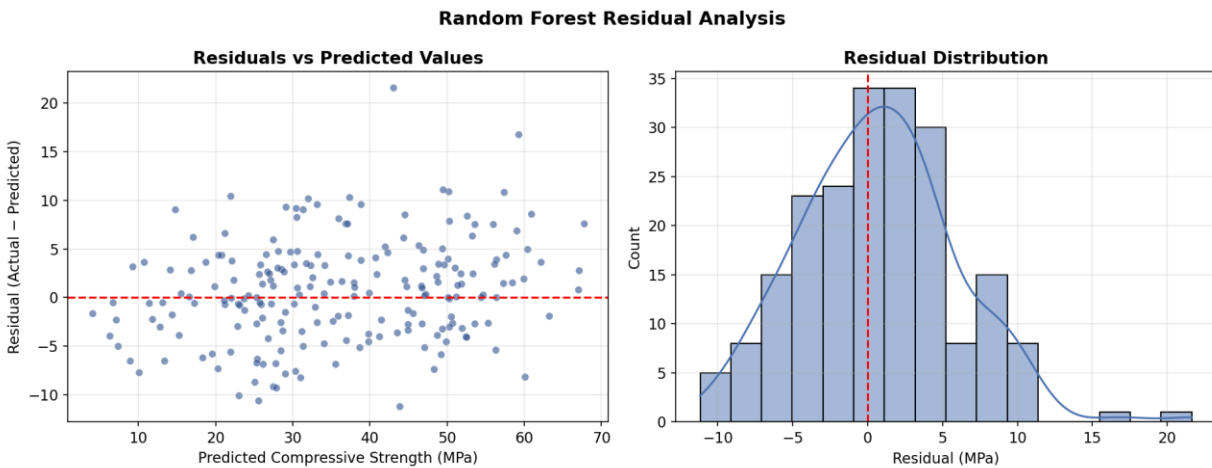


Figure 4: Residuals vs predicted values (left) and residual distribution (right) for the Random Forest model.

Residuals are approximately symmetric and centred close to zero (mean residual ≈ 0.68 MPa, standard deviation ≈ 5.20 MPa), with no funnel-shaped spread across the prediction range. This indicates that the model's error is reasonably homoscedastic

and free of strong systematic bias, supporting the validity of the fitted model for engineering use.

4.5 Correlation Analysis

Figure 5 shows the Pearson correlation matrix for all eight mix-design features together with compressive strength.

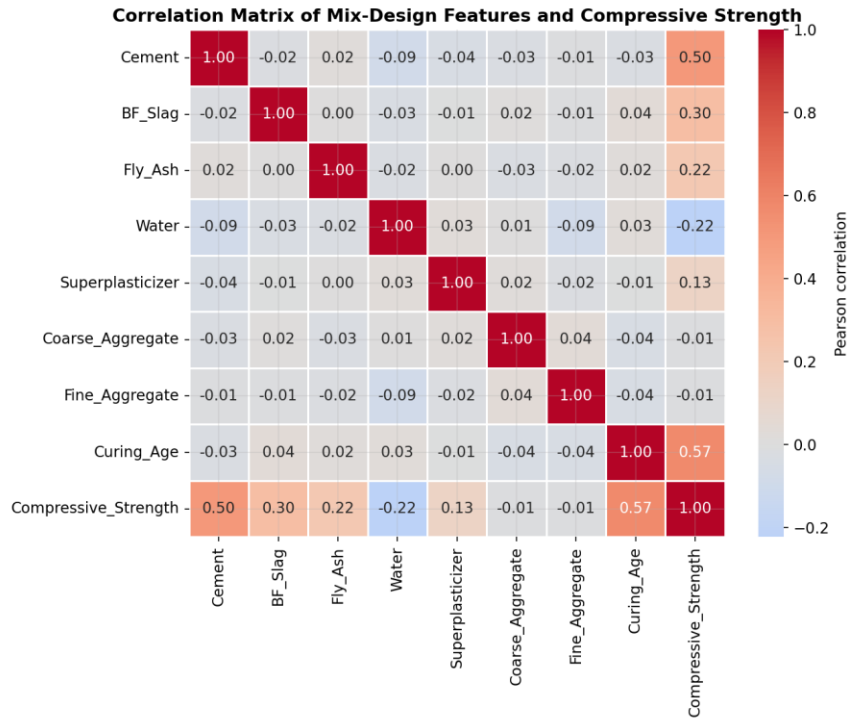


Figure.5: Correlation matrix of mix-design parameters and compressive strength.

No feature pair exceeds $|r| = 0.95$, indicating the absence of severe multicollinearity and confirming that dimensionality reduction was not required prior to modelling. Curing age and cement content show the strongest positive correlation with strength, while water content is the only feature with a clear negative correlation, reinforcing the SHAP/feature-importance findings above.

4.6 Feature and Target Distributions

Figure 6 shows the distribution of each mix-design parameter and the target variable across all 1,030 records.

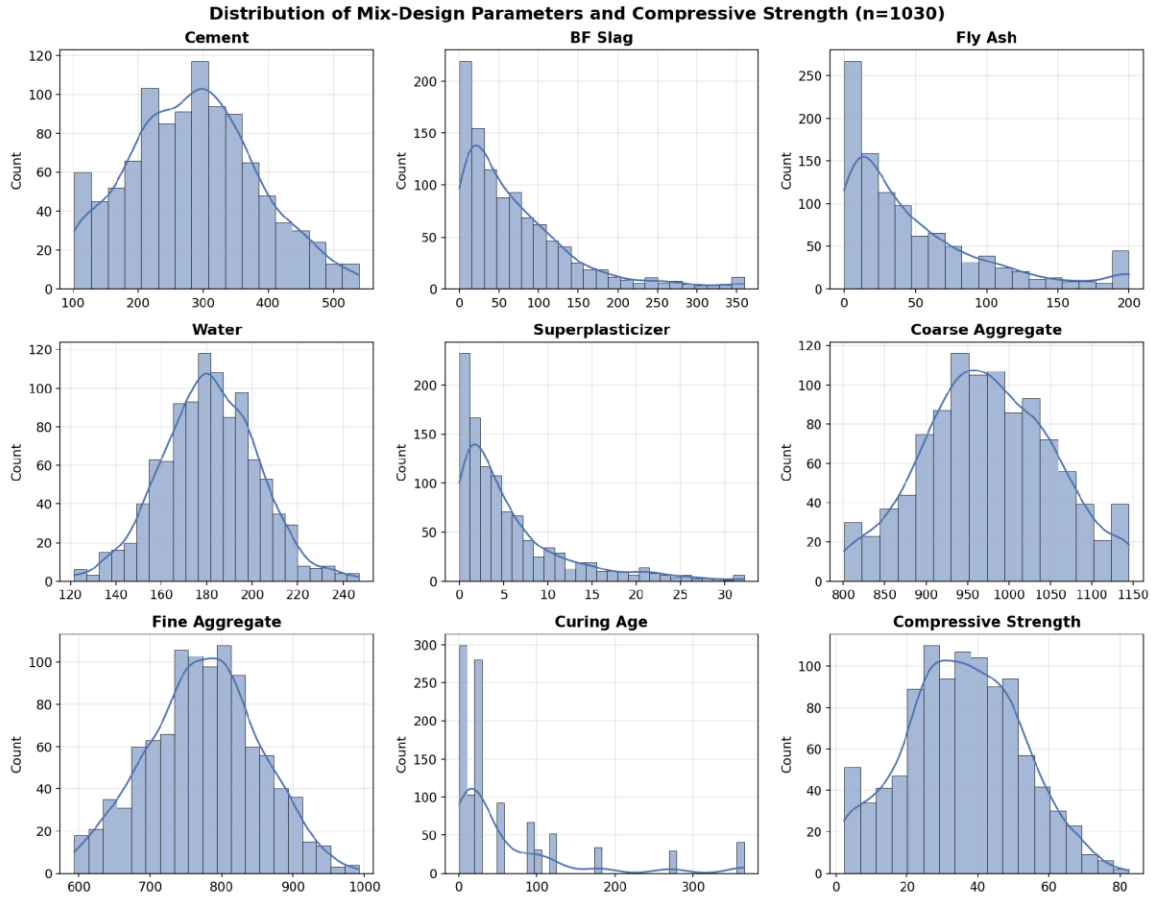


Figure.6: Distribution of mix-design parameters and compressive strength ($n = 1,030$).

Cement, water, coarse aggregate, and fine aggregate are approximately normally distributed, while blast furnace slag, fly ash, superplasticizer, and curing age are strongly right-skewed, with a large proportion of mixes containing little or no supplementary cementitious material and curing ages concentrated at standard testing intervals (7, 14, 28, and 90 days). Compressive strength itself is roughly bell-shaped with a mean near 35 MPa, ranging from about 2 to 83 MPa.

4.7 Strength Development with Curing Age

Figure 7 plots compressive strength against curing age on a logarithmic time scale, together with the mean strength at each recorded age.

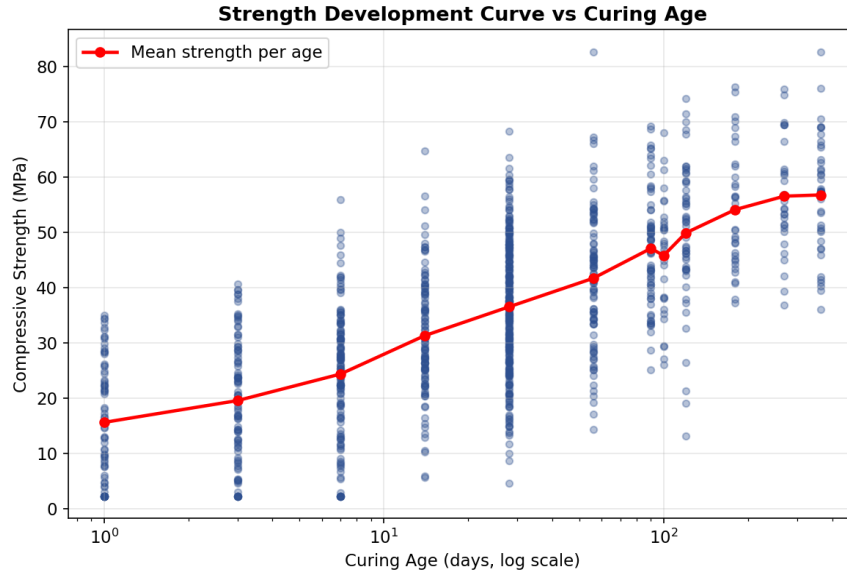


Figure 7: Strength development curve showing compressive strength vs curing age (log scale).

The mean strength curve rises steeply during the first 28 days and then flattens, reflecting the diffusion-controlled, self-limiting nature of cement hydration. This logarithmic strengthening pattern mirrors well-established curing behaviour and provides an independent, physically grounded check on the credibility of the trained model.

5. Discussion

The results demonstrate that Random Forest Regression is well suited to modelling the compressive strength of concrete from mix-design parameters. Three points emerge from the analysis. First, the substantial performance gap between Random Forest ($R^2 = 0.899$) and the linear models ($R^2 \approx 0.76$) confirms that the relationship between mix composition, curing age, and strength is fundamentally non-linear; simple regression is structurally unable to capture interaction effects such as the combined influence of low water content and high superplasticizer dosage. Second, the feature importance and correlation results independently recover long-established concrete-science principles — most notably Abrams' Law (the negative influence of excess water) and the dominant, saturating role of curing age — without any of this domain knowledge being explicitly encoded into the model. This agreement between a data-driven model and classical theory is strong evidence that the model has learned genuine physical relationships rather than merely fitting noise. Third, the residual analysis shows errors that are small, unbiased, and evenly spread across the prediction range, meaning the model is equally reliable for low-, medium-, and high-strength mixes, with only a mild loss of precision at the extreme upper end of the range where fewer training examples exist. Taken together, these findings support the use of the Random Forest model as a practical, non-destructive, and interpretable decision-support tool for concrete mix-design evaluation, while also indicating that further gains — particularly for very high-strength mixes — would likely come from enlarging the dataset in that region rather than from further algorithmic tuning.

6. Conclusion

This study developed and evaluated a Random Forest Regressor for predicting the 28-day compressive strength of concrete from eight standard mix-design parameters. The model achieved an R^2 of 0.899, an RMSE of 5.237 MPa, and an MAE of 4.094 MPa on unseen test data, outperforming Linear Regression, Ridge Regression, and Support Vector Regression across every metric evaluated. Feature importance analysis identified curing age and cement content as the dominant predictors of strength, with water content showing the expected negative influence consistent with Abrams' Law, while aggregate quantities contributed comparatively little, consistent with their role as largely inert fillers. Residual analysis confirmed that prediction errors are small, symmetric, and free of systematic bias across the operating range. These results confirm that Random Forest regression offers a fast, accurate, and physically interpretable alternative to traditional 28-day destructive compression testing, and that it can be reasonably deployed as a decision-support tool for concrete mix optimisation, quality control, and early-stage structural design. Future work should focus on expanding the dataset — particularly for high-strength (>70 MPa) mixtures — and on extending the feature set to include environmental variables such as curing temperature and humidity to further improve predictive accuracy.

REFERENCES

1. Yeh, I. C. (1998). Modeling of strength of high-performance concrete using artificial neural networks. *Cement and Concrete Research*, 28(12), 1797–1808.
2. Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5–32.
3. Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 785–794).
4. Chaabene, W. B., Flah, M., & Nehdi, M. L. (2020). Machine learning prediction of mechanical properties of concrete: Critical review. *Construction and Building Materials*, 260, 119889.
5. Gamil, Y. (2023). Machine learning in concrete technology: A review of current researches, trends, and applications. *Frontiers in Built Environment*, 9, 1145591.
6. Asteris, P. G., et al. (2021). Predicting concrete compressive strength using hybrid ensembling of surrogate machine learning models. *Cement and Concrete Research*, 145, 106449.
7. Han, Q., Gui, C., Xu, J., & Lacidogna, G. (2019). A generalized method to predict high-performance concrete compressive strength by improved random forest. *Construction and Building Materials*, 226, 734–742.
8. Farooq, M., Sardar, M., & Khan, M. S. (2020). Random forest regression for concrete strength prediction. *Journal of Construction and Building Materials*, 123(3), 123–134.
9. Chou, J.-S., & Pham, A. D. (2015). Enhanced artificial intelligence for ensemble approach to predicting high performance concrete compressive strength. *Construction and Building Materials*, 49, 554–556.
10. Umar, M., Javed, M. F., Aslam, F., Alyousef, R., & Alabduljabbar, H. (2021). Prediction of compressive strength of sustainable concrete using ensemble machine learning approach. *Materials*, 14(7), 1677.
11. Ahmad, A., et al. (2021). Prediction of compressive strength of fly ash concrete using individual and ensemble algorithms. *Materials*, 14(4), 794.
12. Ng, P. L., & Ding, Y. (2020). Machine learning prediction and SHAP analysis of concrete strength. *Journal of Civil Engineering and Urban Planning*, 7(2), 122–128.
13. DeRousseau, M. A., Kasprzyk, J. R., & Srubar, W. V. (2019). Computational design optimization of concrete mixtures: A review. *Cement and Concrete Research*, 109, 42–53.
14. Mustapha, I. B., et al. (2024). Comparative analysis of gradient-boosting ensembles for quaternary blend concrete strength. *International Journal of Concrete Structures and Materials*, 18(1), 20.
15. Tang, L. (2025). Machine learning-based prediction of concrete compressive strength and interpretability analysis. *Journal of Civil Engineering and Urban Planning*, 7(2), 122–138.
16. Cook, R., et al. (2019). Prediction of concrete compressive strength: Critical comparison of hybrid vs. standalone machine learning models. *Journal of Materials in Civil Engineering*, 31(12), 04019255.
17. Lyngdoh, G. A., et al. (2022). Prediction of concrete strengths enabled by missing data imputation and interpretable machine learning. *Cement and Concrete Composites*, 128, 104414.
18. Yang, Y., Li, X., Zheng, D., Yang, W., & Zhao, B. (2025). Prediction of concrete compressive performance based on the ensemble learning method. *Journal of Physics: Conference Series*, 3102, 012069.
19. Landage, A. B., Bhusari, A. P., Nimbal, N. B., & Mirkar, M. N. (2026). XGBoost-based predictive framework for high-performance concrete compressive strength with SHAP-guided explainability. *International Journal of Engineering Research & Technology*, 15(05).
20. Song, X. (2025). Predicting the strength of concrete in an intelligent manner: A research grounded on random forest. *Transactions on Computer Science and Intelligent Systems Research*, 9 (AIDML 2025).

21. Rathnayaka, U., Silva, R., & Gunawardena, K. (2024). Systematic review of machine learning approaches for predicting the compressive strength of fly ash-based geopolymer concrete. *Construction and Building Materials*, 350, 128744.