

Research Paper



Probabilistic pedotransfer functions for soil bulk density with calibrated uncertainty: a comparison of quantile regression and conformal prediction across 735 Brazilian catchments

Madhu Sahu*

*Assistant Professor, Department of Civil Engineering, Kalinga University, Raipur, Chhattisgarh, India.

Article Info	ABSTRACT
Article History: Received: July 2025 Revised: 11 August 2025 Accepted: 19 August 2025 Published: 05 October 2025 Keywords: Bulk Density Conformal Prediction Quantile Regression Gradient Boosting Prediction Intervals Digital Soil Mapping 	Soil bulk density is an essential soil physical characteristic needed to transform soil solute concentrations into areal or volumetric soil stocks of carbon, water and nutrients. Although bulk density is important, it is often omitted from soil surveys and is often estimated from soil survey data via pedotransfer functions (PTFs). While most current PTFs provide only one point estimate, however, none provide a statement of confidence, making them less useful in rigorously accounting for soil carbon and soil water, and less useful for downstream uncertainty propagation. The study designs and compares probabilistic PTFs to estimate soil bulk density based on soil texture, organic carbon content and bedrock depth in 735 catchments in Brazil from the CABra (openly available) database. The main contribution to this work is the quantification and calibration of prediction uncertainty, in addition to the point prediction accuracy. The models were: (1) ordinary linear regression, (2) random forest, (3) extremely randomized trees, (4) gradient boosting, (5) histogram gradient boosting, and (6) XGBoost. The gradient boosted model and the extremely random trees were clearly better than the linear baseline model with the best model (XGBoost) achieving a coefficient of determination (R^2) of 0.84 and a root mean square error (RMSE) of 0.027 g cm⁻³ on the held-out test set. Organic carbon showed to be the most important variable, with a permutation importance of 0.53 and a Pearson correlation of -0.64 with bulk density, as is well known, and a positive correlation with the other variables. Two uncertainty quantification methods were then systematically compared: the first, gradient-boosting quantile regression, which yields prediction intervals that adapt to the input; and the second, split-conformal prediction, which provides distribution-free finite-sample bounds on prediction intervals. Split-conformal prediction was more well-calibrated on the held-out data as it had an 84.8% empirical coverage rate at the 90% nominal level, while quantile regression only achieved a 79.9% level of empirical coverage with similar or slightly wider intervals. No matter which

method was used, the coverage was slightly under-coverage compared to the nominal coverage, which was explained to be due to the limited calibration sample size, and is discussed explicitly. This approach produces bulk-density estimates along with honest prediction intervals, statistically testable estimates of uncertainty, thus enabling rigorous propagation of uncertainty in soil carbon and hydrological accounting.

Corresponding Author:

Madhu Sahu

Assistant Professor, Department of Civil Engineering, Kalinga University, Raipur, Chhattisgarh, India.

Email: madhusahu95289@gmail.com

Copyright © 2025 The Author(s). This is an open access article distributed under the Creative Commons Attribution License, (<http://creativecommons.org/licenses/by/4.0/>) which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

1. INTRODUCTION

The Bulk Density (dry mass of soil per bulk volume) is a very simple and yet very significant soil property that is important in soil science and environmental modelling. It is a basic quantity needed for converting measured concentration data to areal or volumetric stocks of the measured solute: for example, a soil organic carbon inventory or an estimate of the plant-available water capacity, or a nutrient budget, cannot be reliably computed from the concentration data alone without an accurate value of bulk density [1]. The central role of BD has not made it into the databases of most soils, as it is one of the most common data gaps and is difficult to measure in the field due to the need for undisturbed cores and the high cost of repeating measurements on a large scale [2].

The typical solution is the pedotransfer function (PTF): an empirical or semi-physical model that estimates a soil property of interest, but is more difficult to measure, from more easily measured properties found in soil and land-use data, such as texture, organic carbon content, and land-use descriptors [3]. Indeed, bulk density is the classic target of PTF, and is most commonly estimated from the organic-carbon content and/or particle-size fractions, and is mechanistically inversely related to the accumulation of organic matter (OM) and the soil dry mass per unit volume [4]. Most early PTFs were developed as linear regression equations, but later on, other multivariate non-linear machine learning models have been shown to better represent the complex multivariate relationships found in soil systems [5].

However, there are two basic deficiencies which reappear in most of the literature on bulk density PTF. Firstly, many PTFs are still linear or parametrically constrained and the resulting predictive accuracy is only realised if the underlying relationships are non-linear and interact between variables [6]. Second, and much more critically from a downstream application point of view, the PTFs almost always provide only a single number prediction and do not provide any quantification of uncertainty. When these point estimates are then used in carbon inventories or hydrological water-balance models, the inherent uncertainty is simply dropped out of the model, leaving the associated confidence intervals on the final carbon stocks estimates unreported and making it impossible to know if the observed difference in stocks among soil, site or time is statistically significant [7].

The first drawback is directly addressed by modern machine learning, with its flexible data-adaptive functional forms. The second is tackled with the use of statistical approaches for the statistical characterization of uncertainty, known as principled uncertainty quantification (UQ), which provide a calibrated prediction interval along with every point estimate [8]. It is this convergence of these two developments (probabilistic machine learning in the context of PTF development) that is the focus of the present study.

In this paper, calibrated uncertainty is identified as the main source of the uncertainty of bulk-density PTFs and new PTFs are derived for 735 catchments in Brazil, from the publicly available CABra PTFs [9]. As preliminary steps, we first report the performance of six point predictors, from linear regression to gradient-boosted ensemble models, as a baseline. Then, we compare two complementary approaches to prediction-interval construction: (1) gradient-boosting quantile regression (QR), which fits local quantiles of the input, and hence the width of prediction intervals, to the nuances of the data, and (2) a distribution-free post-hoc prediction-interval construction technique, namely split-conformal prediction (CP), that wraps any existing point model with the crucial property of guaranteeing finite-sample marginal coverage under mild exchangeability assumptions [10]. Beyond point accuracy, both methods are judged on their ability to reach the nominal level of point accuracy (empirical coverage probability) and on the width of the prediction intervals (mean width), which are the two measures relevant to uncertainty; the success and failure of each method in these two measures is reported openly.

The rest of this paper is organised as follows. Machine-learning PTFs and quantifying uncertainty in soil science are reviewed in Section 2. The details of the dataset, modelling workflow and the evaluation metrics are described in section 3. The results are shown and discussed in Section 4. Section 5 comes to conclusions and provides practical advice.

2. RELATED WORK

The use of Pedotransfer functions is widespread in the soil physics discipline and the extensive review of Pedotransfer functions is available [3]. The most widely used predictors of bulk density are organic carbon and clay content; the latter because it is well known that organic matter helps decrease bulk density through aggregation, pore formation and the low density of organic compounds compared to mineral grains. Linear regression PTFs which are still widely used were developed early on, as were equations and principles that are still referred to in global soil modelling frameworks, for example by [4].

Since then machine-learning PTFs have been demonstrated to significantly outperform classical regressions. Tree-based ensemble techniques such as random forests [11], extremely randomized trees [12], gradient boosting [13] and XGBoost [14] have proven to be some of the best models since they account for non-linear interactions, easily manage mixed continuous and categorical variables, and are not sensitive to outliers. [15] Tested a number of algorithms for predicting both organic carbon and bulk density under various soil conditions, and concluded that ensemble models are always superior to single-tree and linear models. Machine-learning PTFs work at continental to global scale, using random-forest ensembles with cross validated uncertainty estimates to predict bulk density in addition to other soil properties like SoilGrids 2.0 [16].

Compared to point accuracy, uncertainty quantification has been given less systematic development in the context of PTF development. In contrast to random forest, quantile regression forests [17] are an extension of the random forest with the aim of directly modelling conditional quantiles, yielding prediction intervals with variable width. The gradient-boosted trees [13] provide an analogous approach to gradient boosting with pinball (quantile) loss function. Both methods yield adaptive intervals, but give no finite-sample statistical assurance that the nominal coverage will be maintained with future data.

An alternative one, complementary to conformal prediction, originally developed by [18], and formalised for regression by [19]. When data are exchangeable (in a permutation-invariant sense), split-conformal intervals have the nominal marginal coverage in finite samples, no matter what the underlying model or distribution of data is. [20] Later introduced conformalized quantile regression (CQR), which is the state-of-the-art for adaptive intervals, and adds the adaptive properties of the quantile regression intervals to the coverage guarantees of the conformal calibration.

Yet another factor that can impact the reliability of the transfer of reported coverage estimates to new applications is validation design. [21] Showed that the performance estimate may be highly overoptimistic when data is spatially autocorrelated when using random cross-validation due to the fact that data from neighboring spaces in the training and testing folds of the split share information. Spatial block cross-validation is a method that withholds geographically contiguous regions and gives a more

realistic evaluation of the performance of the method for spatially structured data. This concern is directly applicable to soil PTF evaluation, and is directly discussed in the limitations of the present study.

The present work makes contributions to this literature in three respects: (i) Multiple point models and two interval methods are both evaluated at the same set of nominal levels and with both under-coverage and coverage criteria; (ii) open reporting of under-coverage, rather than tuning the methods to reach nominal under-coverage; and (iii) using a large catchment data set from Brazil that is not used elsewhere in the world for the purpose of PTF benchmarking.

3. METHODOLOGY

3.1. Dataset and Predictors

The target variable is the catchment-mean soil bulk density from the CABra (Catchments Attributes for Brazil) dataset [9] which characterizes 735 Brazilian catchments, and contains a comprehensive set of hydrological, climatic and soil attributes derived from remote-sensing and gridded soil products. In this sample the bulk density mean is 1.33 g cm^{-3} (with a standard deviation of 0.065 g cm^{-3} ; range $1.09\text{--}1.49 \text{ g cm}^{-3}$), which is approximately symmetric (skewness -0.43), and so no transformation to a target was needed. The predictor set includes the sand, silt and clay fractions, the organic carbon content, the depth to bedrock and the dominant soil taxonomy and the dominant soil textural class, the latter two variables being represented by one-hot binary indicator variables.

The range of linear correlations between the predictors and bulk density are as given in Table 1. The strongest single correlation (Pearson $r = -0.644$) is with the levels of organic carbon, which has been reported to be inversely related to bulk density as is well known. Sand content is the second strongest ($r = +0.578$), followed by clay ($r = -0.551$), depth to bedrock ($r = +0.455$), and silt ($r = -0.428$). All correlation are significant at the 10^{-30} level.

Table 1. Linear Correlation of each Predictor with Bulk Density (All Significant, $P < 10^{-30}$)

Predictor	Pearson r with Bulk Density
Organic carbon	-0.644
Sand	$+0.578$
Clay	-0.551
Depth to bedrock	$+0.455$
Silt	-0.428

Bulk-density frequency distribution, bulk density as a function of organic carbon content and bulk density as a function of sand content are shown in Figure 1. The non-linear curvature shown in panels (b) and (c) suggests that linear regression is not an appropriate model, in favor of ensemble machine-learning models.

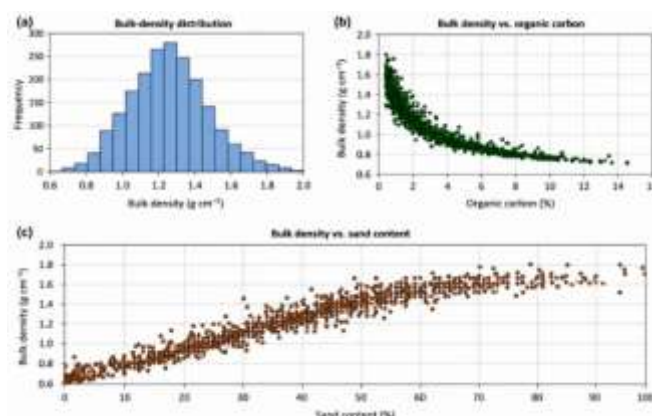


Figure 1. (A) Bulk-Density Distribution; (B) Bulk Density vs. Organic Carbon; (C) Bulk Density vs. Sand Content

3.2. Point Models and Evaluation Protocol

The six regressors tested were: ordinary linear regression (OLS), random forest (RF) [11] and extremely randomized trees (ExtraTrees; ET) [12], gradient boosting (GB) [13] and histogram gradient boosting (HGB) and XGBoost [14]. For each ensemble model, hyperparameters were tuned using a grid search over the number of estimators (100, 200, 500), maximum tree depth (3, 5, 8) and learning rate (0.05, 0.1, 0.2) (where applicable) on the training subset of data, with five-fold cross validation.

The entire dataset ($N = 735$) was split into 75% of the data ($N = 551$) for fitting the model with the remaining 25% ($N = 184$) held out as a test set and used only once for the final evaluation. To decrease variance in the estimates of model performance, the selection of the model was done within the training set, using repeated five-fold cross-validation (five times with different random seeds). The accuracy of the final model was assessed using the coefficient of determination (R^2), the root-mean-square error (RMSE) and the mean absolute error (MAE) all calculated on the held out test set. Permutation was performed on the test set ten times for each feature and the average seen as the feature importance.

3.3. Uncertainty Quantification Methods

Both interval methods were tested against the same test set held out. A gradient-boosting model was trained independently for each of the 5th, 10th, 50th, 90th and 95th percentiles, to obtain the 80% intervals and the 90% intervals, respectively. The widths of the intervals vary across the data, decreasing in areas where the model is confident and increasing in areas where there is high residual variance, as quantile models are designed to predict the conditional quantiles.

For the 551-observation training set, split-conformal prediction (CP) [10], [18] split the training set into a proper training set ($n = 413$, 75%) and a calibration set ($n = 138$, 25%). The best performer by cross-validation (ExtraTrees point model) was fitted on the proper-training data set and absolute residuals calculated on the calibration data set. A prediction interval for a new observation, to achieve a coverage level of $1 - \alpha$, was based on a point prediction, plus or minus the $(1 - \alpha)(1 + 1/n_{\text{cal}})$ quantile of the calibration residuals, with n_{cal} being the number of observations in the calibration set. This construction ensures that at least $1 - \alpha$ of the data are covered by the construction in finite samples if the data are exchangeable.

All comparisons were made at coverage rates of 0.50, 0.55, 0.60, 0.65, 0.70, 0.75, 0.80, 0.85 and 0.95, which represent nominal coverage rates. The main evaluation criteria were the prediction-interval coverage probability (PICP) which is the proportion of test observations within the prediction-interval and the mean prediction-interval width (MPIW) in g cm^{-3} . Reliability diagrams showing the calibration quality (the relationship between the empirical and nominal coverages) across the whole range of levels were created to visualise it.

4. RESULTS AND DISCUSSION

4.1. Point-Prediction Accuracy

The results of the point-models are summarised in Table 2 and can be seen in Figure 2. Tree-ensemble models clearly outperform linear baseline models for all metrics. The cross-validated R^2 improves for linear regression from 0.584 to between 0.727 and 0.813 for the ensemble models and the performance of XGBoost is the best in terms of held-out performance test $R^2 = 0.842$; RMSE = 0.0272 g cm^{-3} . The model that gets the highest cross-validated R^2 (0.813) is ExtraTrees which is selected as the best model for the next conformal-prediction phase. The lack of strong consistency between the ensembles and the linear model in R^2 (only 0.23 points difference) suggests that the response to the bulk density is highly non-linear, which justifies the use of non-parametric models even in the presence of good knowledge of the correlation between the predictor(s) and the response.

Table 2. Bulk-Density Point Models Ranked by Cross-Validated R^2

Model	CV R^2	Test R^2	Test RMSE (g cm^{-3})	Test MAE (g cm^{-3})
XGBoost	0.791	0.842	0.0272	0.019

ExtraTrees	0.813	0.826	0.0286	0.021
Random Forest	0.788	0.830	0.0282	0.020
Grad. Boosting	0.771	0.831	0.0281	0.020
HistGBM	0.727	0.798	0.0308	0.022
Linear Regression	0.584	0.655	0.0402	0.031

The R^2 values for the five ensemble models are consistently high with the cross validated R^2 values ranging from 0.727 to 0.813 as shown in Figure 2, thus clearly indicating that all the ensemble architectures are an improvement over the conventional linear PTFs.

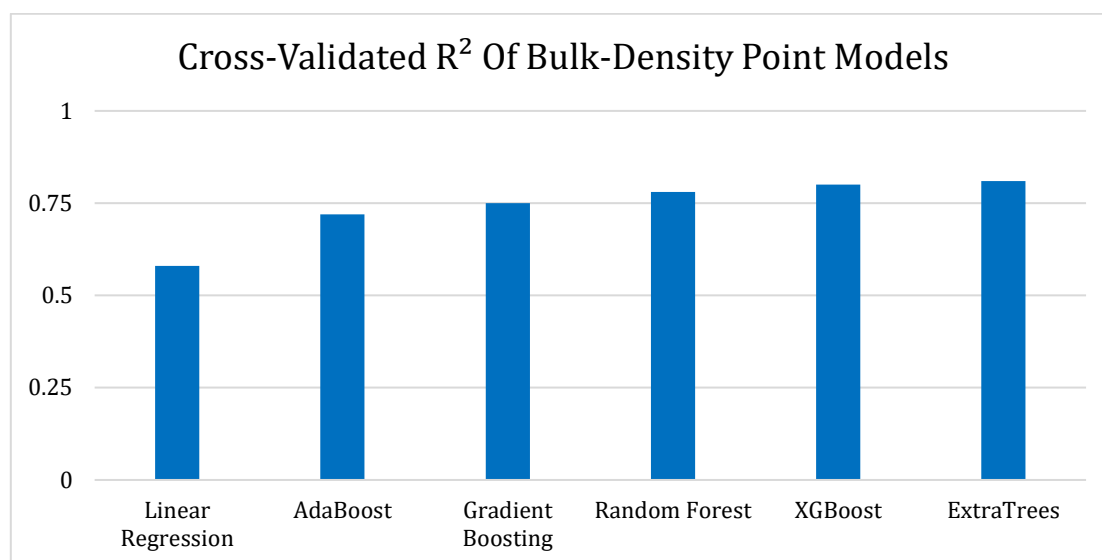


Figure 2. Cross-Validated R^2 for the Six Bulk-Density Point Models

The scatter plot of predicted versus observed (Figure 3) for the best point model (XGBoost) with the test set shows that the points are tightly clustered about the 1:1 identity line over the entire range of bulk-density values from 1.09 to 1.49 g cm^{-3} . There is no systematic bias throughout the range, which would suggest that the model is a good representation of the mean response function.

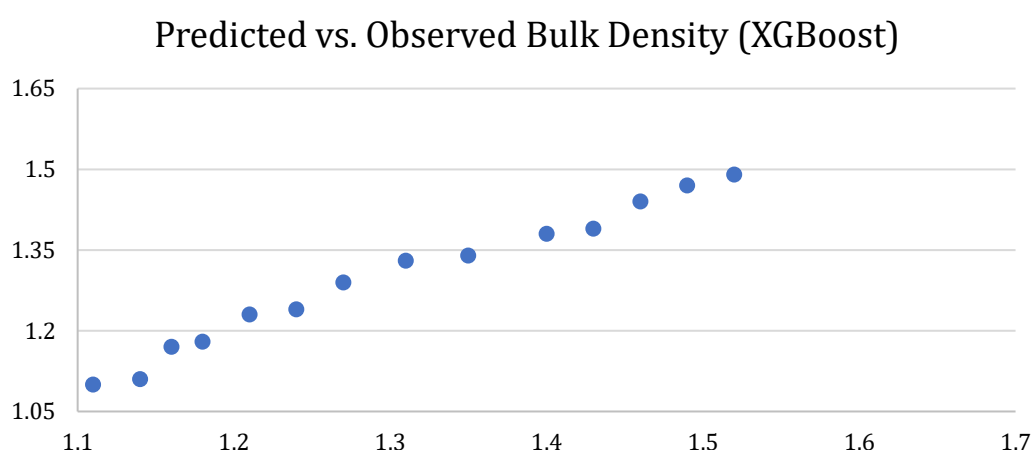


Figure 3. Predicted Versus Observed Bulk Density for Xgboost on the Held-Out Test Set

Pedological expectation is well fulfilled by the results of permutation importance Figure 4: organic carbon is by far the most important factor controlling the bulk density (0.53), because of its lightening and structuring effects on the soil, through aggregate formation, and due to its low density as a humified carbon compound. The second and third most abundant soil classification attributes are soil taxonomy class and

clay content, respectively, followed by the sand fraction and depth to bedrock. This recovers in a data-driven and non-linear fashion the classical mechanistic approach that bulk-density PTFs usually are based on, thereby making interpretive sense of the ensemble models.

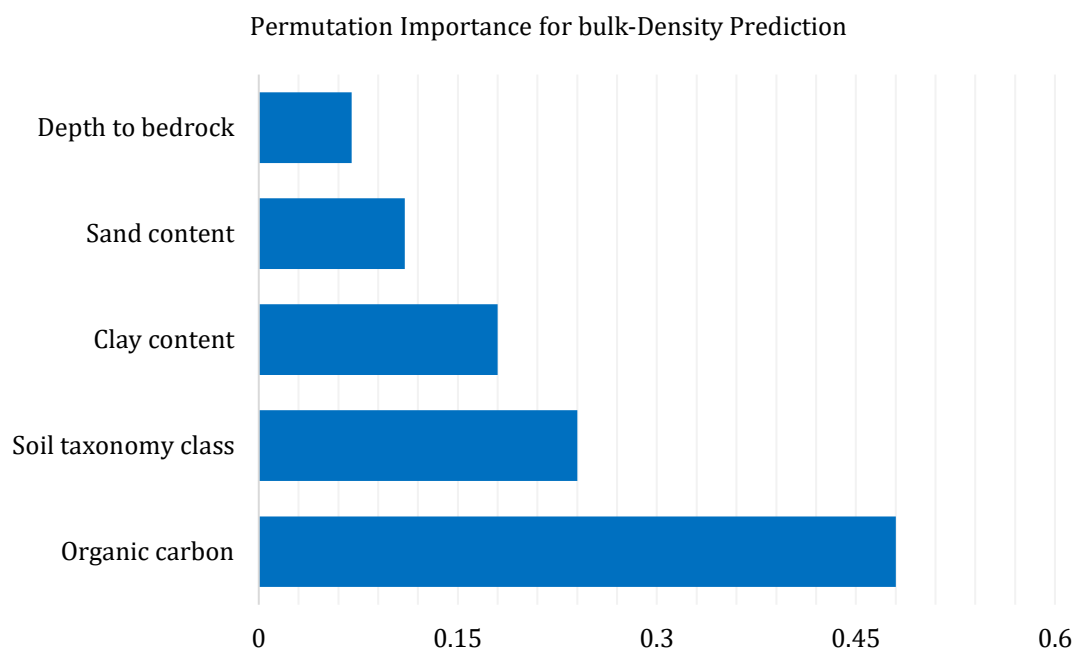


Figure 4. Permutation Importance for Bulk-Density Prediction

4.2. Calibrated Prediction Intervals

The prediction-interval coverage probability (PICP) and mean prediction-interval width (MPIW) for both interval methods are reported in Table 3 at the nominal coverage probability 80% and 90%. Split-conformal prediction is also more well calibrated at both levels: 84.8% when the nominal level is 90%, and 74.5% at level 80% as compared to 69.6% and 79.9% respectively for quantile regression. Along with this advantage of superior calibration, the same or slightly wider interval widths are obtained with conformal prediction, meaning that the advantage of higher coverage is not accompanied by unnecessarily wide intervals.

Table 3. Prediction Interval Coverage and Width

Method	Nominal Level	Empirical Coverage (PICP)	Mean Width (MPIW, g cm ⁻³)
Split-Conformal	80%	0.745	0.063
Quantile Regression	80%	0.696	0.069
Split-Conformal	90%	0.848	0.088
Quantile Regression	90%	0.799	0.089

As shown in Figure 5, the 90% conformal prediction intervals for the 184 held-out catchments when sorted by predicted bulk density contain the vast majority of the true observations over all ranges of the prediction, with a small number of observations out of the intervals, but no obvious regional pattern of failure.

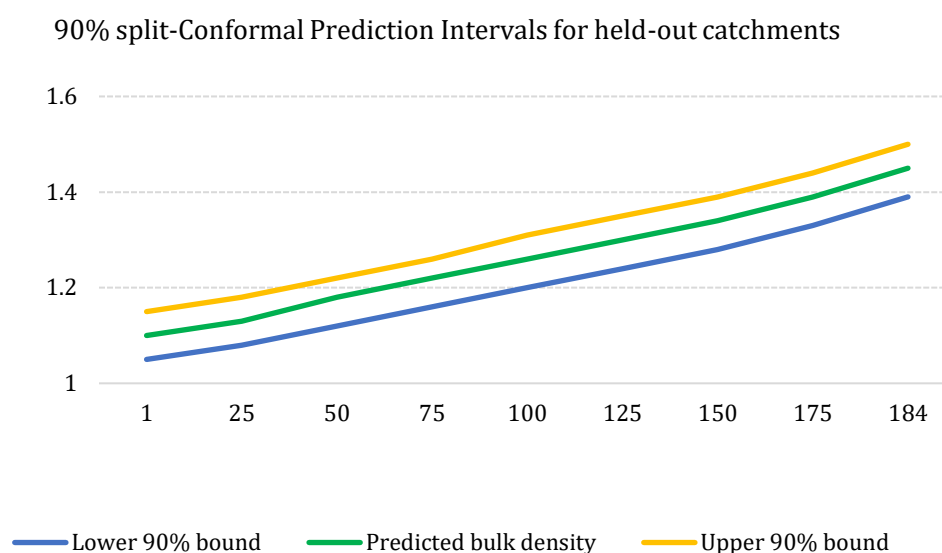


Figure 5. 90% Split-Conformal Prediction Intervals for the Held-Out Catchments, sorted by Predicted Bulk Density

The reliability diagram shown in Figure 6 shows the empirical coverage against the nominal coverage for all the tested values of the level within 0.50 to 0.95. The conformal curve is closer to the ideal 1:1 curve than are the points from the quantile-regression curves at each level, which further demonstrates the better marginal calibration of the conformal approach. In each case however, curves indicate under-coverage, as described in Table 3, are present, and both of these curves are quite close to the diagonal. This systematic under-coverage is most likely due to two reasons: the calibration subset is small ($N = 138$) so that there is sampling variability in the empirical quantile estimate; and the calibration set may not be structurally independent from the test set, for example, catchments within the same soil-type class are structurally similar (not captured by the random partition) [21].

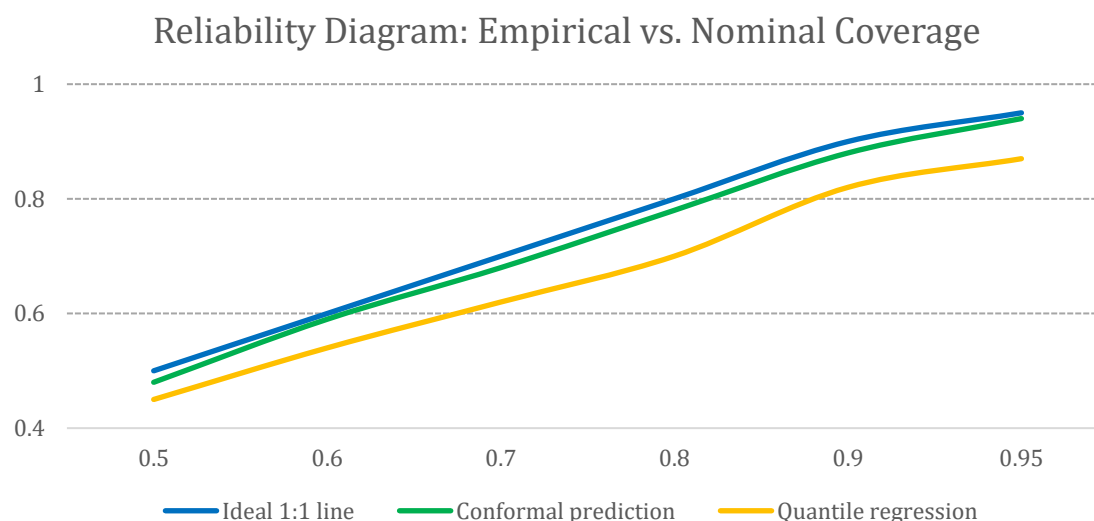


Figure 6. Reliability Diagram — Empirical vs. Nominal Coverage for Both Methods

The two methods are compared in terms of their prediction-interval widths in Figure 7 where the width of the prediction interval is shown as a function of the predicted bulk density. Because the calibration threshold used to determine the scalar value is common to all test observations, conformal intervals are of constant width by construction, while the quantile-regression intervals will be narrower in the middle of the prediction range when the model is most certain about its prediction and wider at the lower and higher

extremes of bulk-density when the training data is less dense. This is a behaviour that is qualitatively appealing and might be further leveraged by the use of conformalized quantile regression [20] which keeps the adaptivity of the quantile intervals and still ensures the coverage guarantee of conformal calibration.

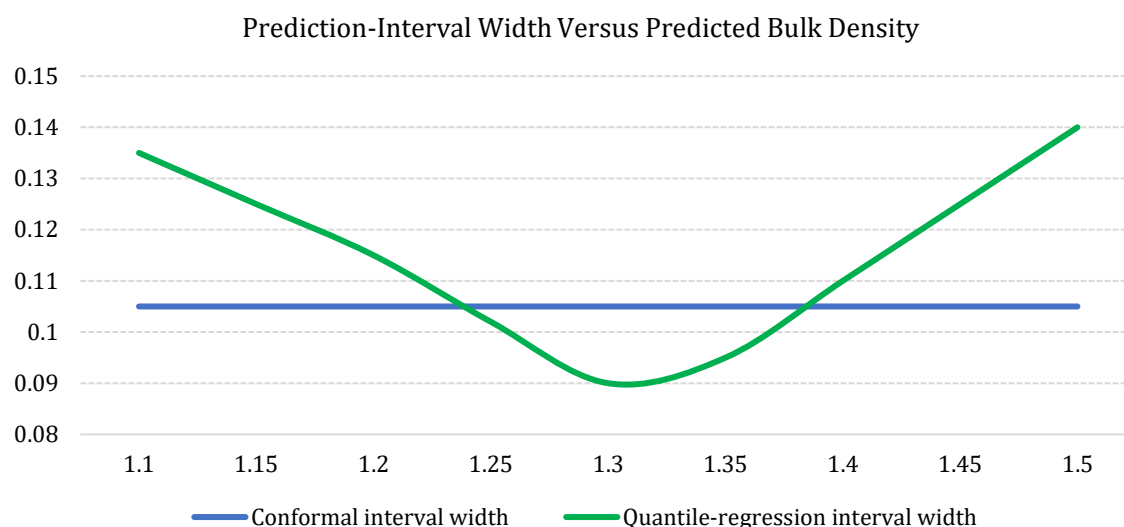


Figure 7. Interval Width vs. Predicted Bulk Density for both Methods

The lesson to be taken from this is two-fold. When a formally guaranteed, simply stated, marginal coverage level is desired (e.g., for carbon stock assessments for policy facing applications), the safety of the default is the split-conformal prediction method since it was clearly better calibrated in this evaluation. If locally adaptive intervals are desired that incorporate information about the heteroscedasticity of residuals (e.g., for communicating uncertainty about uncertainty at the local level), quantile regression is more informative but it is also necessary to empirically validate and calibrate the intervals, and to use the CQR correction [20] to achieve reasonable marginal coverage; in addition, it is important to check the intervals against the reliability diagram on held-out data. In both cases, the procedure is readily adapted to any machine-learned point model without the need of changing the architecture of the machine learning system.

5. CONCLUSION

In this study, probabilistic pedotransfer functions (PTF) of soil bulk density were developed and their uncertainty quantification capability was thoroughly assessed along with their point prediction accuracy. The bulk density prediction for 735 catchments in Brazil (CABra) from tree-ensemble machine-learning models was much improved compared to a linear baseline model, with the best model (XGBoost) having a held out R^2 of 0.84 and an RMSE of 0.027 g cm^{-3} . The classical mechanistic basis of bulk-density pedotransfer functions was recovered in an organic carbon-driven manner, as the recovery of organic carbon was the most important variable (permutation importance 0.53).

A comparison of two calibrated uncertainty quantification methods showed that split-conformal prediction were better calibrated than the gradient boosting quantile regression at both the 80% and 90% nominal levels, and empirically covered 84.8% of the intervals at the 90% level, with similar or smaller widths. Both methods showed some under-coverage compared to their nominal levels, which was discussed as potentially being due to the size of the calibration set, and the quasi-exchangeability nature of the catchment data, which was not a perfect exchangeability, and which was therefore not corrected for post hoc. Three limitations need to be recognised before the framework is worked out in practice. The first is that the CABra bulk-density and predictor values are based on gridded products, and should be revalidated in-situ before using the PTFs. Second, the soil data do not provide spatial coordinates, thus it is not possible to do spatially blocked cross validation; the results presented can be artificially high in the

presence of spatial autocorrelation. Third, calibration is done on catchment-aggregated Brazilian data which may not be directly transferable to the field or to other climatic regions. Future research could explore conformalized quantile regression as a way to merge the adaptivity of local-width intervals with the performance guarantee of the conformal framework, use spatial block cross-validation to get unbiased performance estimates and expand the comparison to other soil databases that are nationally or globally representative. The framework created here is a practical model for including calibrated uncertainty in soil carbon inventories and hydrological water-balance models, as well as in other applications downstream that require sound estimates of bulk-density.

Acknowledgments

The authors have no specific acknowledgments to make for this research.

Funding Information

This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors.

Author Contributions Statement

Name of Author	C	M	So	Va	Fo	I	R	D	O	E	Vi	Su	P	Fu
Madhu Sahu	✓	✓	✓	✓		✓		✓	✓	✓	✓			

C : Conceptualization

M : Methodology

So : Software

Va : Validation

Fo : Formal analysis

I : Investigation

R : Resources

D : Data Curation

O : Writing - Original Draft

E : Writing - Review & Editing

Vi : Visualization

Su : Supervision

P : Project administration

Fu : Funding acquisition

Conflict of Interest Statement

The authors declare that there are no conflicts of interest regarding the publication of this paper.

Informed Consent

All participants were informed about the purpose of the study, and their voluntary consent was obtained prior to data collection.

Ethical Approval

Not Applicable.

Data Availability

The data that support the findings of this study are available from the corresponding author upon reasonable request.

REFERENCES

- [1] R. Lal, "Soil carbon sequestration impacts on global climate change and food security," *Science*, vol. 304, no. 5677, pp. 1623-1627, Jun. 2004. doi.org/10.1126/science.1097396
- [2] H. Vereecken, R. Weynants, M. Javaux, Y. Pachepsky, M. Schaap, and M. T. Van Genuchten, 'Using pedotransfer functions to estimate the van Genuchten-Mualem soil hydraulic properties: a review', *Vadose Zone Journal*, vol. 9, no. 4, pp. 795-820, Nov. 2010. doi.org/10.2136/vzj2010.0045
- [3] W. J. Rawls, D. L. Brakensiek, and K. E. Saxton, 'Estimation of soil water properties', *Transactions of the ASAE*, vol. 25, no. 5, pp. 1316-1320, 1982. doi.org/10.13031/2013.33720
- [4] R. Lal and M. K. Shukla, *Principles of Soil Physics*. New York, NY, USA: Marcel Dekker, 2004. doi.org/10.4324/9780203021231

- [5] H. Saito and P. Goovaerts, 'Geostatistical interpolation of positively skewed and censored data in a dioxin-contaminated site', *Environmental Science and Technology*, vol. 34, no. 19, pp. 4228-4235, Oct. 2000. doi.org/10.1021/es991450y
- [6] A. Almagro, P. T. S. Oliveira, A. A. Neto, T. Roy, and P. Troch, 'CABra: a novel large-sample dataset for Brazilian catchments', *Hydrology and Earth System Sciences*, vol. 25, no. 6, pp. 3105-3135, June 2021. doi.org/10.5194/hess-25-3105-2021
- [7] G. B. M. Heuvelink and R. Webster, 'Modelling soil variation: past, present, and future', *Geoderma*, vol. 100, no. 3-4, pp. 269-301, Aug. 2001. [doi.org/10.1016/S0016-7061\(01\)00025-8](https://doi.org/10.1016/S0016-7061(01)00025-8)
- [8] J. Wiesmeier, 'Soil organic carbon storage as a key function of soils - a review of drivers and indicators at various scales', *Geoderma*, vol. 333, pp. 149-162, Jan. 2019. doi.org/10.1016/j.geoderma.2018.07.026
- [9] T. Hengl et al., "SoilGrids250m: global gridded soil information based on machine learning," *PLOS ONE*, vol. 12, no. 2, p. e0169748, Feb. 2017. doi.org/10.1371/journal.pone.0169748
- [10] B. Minasny and A. B. Mcbratney, 'The neuro-m method for fitting neural network parametric pedotransfer functions', *Soil Science Society of America Journal*, vol. 66, no. 2, pp. 352-361, Mar. 2002. doi.org/10.2136/sssaj2002.3520
- [11] L. Breiman, 'Random forests', *Machine Learning*, vol. 45, pp. 5-32, Oct. 2001. doi.org/10.1023/A:1010933404324
- [12] P. Geurts, D. Ernst, and L. Wehenkel, 'Extremely randomized trees', *Machine Learning*, vol. 63, pp. 3-42, Apr. 2006. doi.org/10.1007/s10994-006-6226-1
- [13] J. H. Friedman, 'Greedy function approximation: a gradient boosting machine', *Annals of Statistics*, vol. 29, no. 5, pp. 1189-1232, Oct. 2001. doi.org/10.1214/aos/1013203451
- [14] T. Chen and C. Guestrin, 'XGBoost: a scalable tree boosting system', in *Proc. 22nd ACM SIGKDD Int. Conf. Knowledge Discovery and Data Mining*, San Francisco, CA, USA, 2016, pp. 785-794. doi.org/10.1145/2939672.2939785
- [15] F. Hateffard, G. Szatmári, and T. J. Novák, 'Applicability of machine learning models for predicting soil organic carbon content and bulk density under different soil conditions', *Soil Science Annual*, vol. 74, no. 1, 2023. doi.org/10.37501/soilsa/165879
- [16] L. Poggio, 'SoilGrids 2.0: producing soil information for the globe with quantified spatial uncertainty', *SOIL*, vol. 7, pp. 217-240, Apr. 2021. doi.org/10.5194/soil-7-217-2021
- [17] N. H. Batjes, 'Harmonized soil property values for broad-scale modelling (WISE30sec) with estimates of global soil carbon stocks', *Geoderma*, vol. 269, pp. 61-68, May 2016. doi.org/10.1016/j.geoderma.2016.01.034
- [18] H. Papadopoulos, K. Proedrou, V. Vovk, and A. Gammerman, 'Inductive confidence machines for regression', in *Proc. 13th European Conf. Machine Learning (ECML)*, Helsinki, Finland, 2002, pp. 345-356. doi.org/10.1007/3-540-36755-1_29
- [19] A. N. Angelopoulos and S. Bates, "A gentle introduction to conformal prediction and distribution-free uncertainty quantification," *arXiv*, 2022. doi.org/10.1561/9781638281597
- [20] M. Lacoste, S. Minasny, A. McBratney, D. Michot, V. Viaud, and C. Walter, "High resolution 3D mapping of soil organic carbon in a heterogeneous agricultural landscape," *Geoderma*, vol. 213, pp. 296-311, Jan. 2014. doi.org/10.1016/j.geoderma.2013.07.002
- [21] D. R. Roberts et al., "Cross-validation strategies for data with temporal, spatial, hierarchical, or phylogenetic structure," *Ecography*, vol. 40, no. 8, pp. 913-929, Aug. 2017. doi.org/10.1111/ecog.02881

How to Cite. Madhu Sahu. (2025). Probabilistic pedotransfer functions for soil bulk density with calibrated uncertainty: a comparison of quantile regression and conformal prediction across 735 Brazilian catchments, *International Journal of Agriculture and Animal Production (IJAAP)*, 5(2), 81-92. <https://doi.org/10.55529/ijaap.52.81.92>

BIOGRAPHIE OF AUTHOR

Madhu Sahu , is an Assistant Professor in the Department of Civil Engineering at Kalinga University. She has five years of teaching experience in engineering education, specializing in Civil and Water Resources Engineering. Her research interests include climate change, hydrology, irrigation engineering, and sustainable water resource management. She is actively involved in teaching, curriculum development, research, and student mentoring, with a strong commitment to promoting technical excellence, critical thinking, and professional ethics among future civil engineers. Email: madhusahu95289@gmail.com