

Estimating Reservoir Sedimentation Using Machine Learning

Amanda L. Cox, M.ASCE¹; Deanna Meyer²; Alejandra Botero-Acosta³; Vasit Sagan⁴; Ibrahim Demir⁵; Marian Muste⁶; Paul Boyd, M.ASCE⁷; and Chandra Pathak, F.ASCE⁸

Abstract: Several reservoirs across the United States are filling with sediment, which jeopardizes their functionality and increases maintenance costs. USACE developed the Reservoir Sedimentation Information (RSI) system to assess reservoir aggradation and track dam operation suitability for water resource management and dam safety. The RSI data set contains historical elevation-capacity data for approximately 400 dams (excluding navigation structures), which correspond to less than 1% of dams across the United States. Thus, there is a critical need to develop methods for estimating reservoir sedimentation for unmonitored sites. The goal of this project was to create a generalized method for estimating reservoir sedimentation rates using reservoir design information and watershed data. To meet this objective, geospatial tools were used to build a refined composite data set to complement the RSI system's data with precipitation and watershed characteristics. Nine deep learning models were then used on the benchmark data set to determine its accuracy at predicting capacity loss for the RSI reservoirs: four supervised machine learning models, four deep neural network (DNN) models, and a multilinear power regression model. A DNN model, containing a progressively increasing node and layer construction, was deemed the most accurate, with R^2 values from its calibration and validation data sets being 0.83 and 0.70, respectively. The best model was recalibrated over the entire data set, which showed greater accuracy on the prediction of the RSI reservoir's capacity loss, with an R^2 of 0.81. This predictive model could be used to evaluate the capacity loss of unmonitored reservoirs, forecast sedimentation rates under future climate conditions, and identify reservoirs with the highest risk of losing functionality. **DOI:** 10.1061/JHYEFF.HEENG-6135. © 2024 American Society of Civil Engineers.

Practical Applications: Many communities depend on reservoirs for a variety of socioeconomic benefits, such as providing reliable water sources and flood mitigation. Rivers entering reservoirs are a constant source of silt, sand, and gravel particles (i.e., sediments) that deposit slowly, filling the reservoirs over time, thus reducing their volume capacity and effectiveness. Surveys to measure reservoir capacities are labor intensive and expensive and many sites go unmonitored. This study provides prediction tools to estimate capacity loss over time using data derived from publicly available data sources (e.g., digital elevation models, monthly precipitation data, and the National Inventory of Dams). A key role of this tool is to forecast reservoir capacity loss over time under varying climate change conditions and relate the associated sedimentation processes to local and regional conditions. The tool can also be applied broadly to hindcast capacity loss for reservoirs with or without prior surveys for identifying high-risk sites that should be investigated further. USACE plans to use these prediction tools with the RSI database to conduct a national assessment of reservoir impacts, which will inform distributions of federal resources to address water security concerns related to reservoir sedimentation.

Author keywords: Reservoir sedimentation; Reservoir capacity; Machine learning.

Introduction

Dams and their associated reservoirs enable water storage, flood control, and hydroelectric power generation, and supply reliable water resources for various societal needs. However, reservoirs throughout the nation are slowly filling with sediment, diminishing

their life cycle and reducing their effectiveness, while increasing their cost of maintenance (Sholtes et al. 2018). The immediate consequences of sediment retention in reservoirs are diminishing reservoir capacity, creation of backwater flooding upstream, as well as impairing turbines of the structure (Morris and Fan 1998). The costs of remediating accumulated sediment in these structures may be

¹Associate Professor, WATER Institute, Saint Louis Univ., St. Louis, MO 63103 (corresponding author). ORCID: <https://orcid.org/0000-0003-0719-2669>. Email: amanda.cox@slu.edu

²Graduate Research Assistant, WATER Institute, Saint Louis Univ., St. Louis, MO 63103. Email: deanna.meyer@slu.edu

³Research Scientist, WATER Institute, Saint Louis Univ., St. Louis, MO 63103. ORCID: <https://orcid.org/0000-0002-8458-0326>. Email: alejandra.boteroacosta@slu.edu

Note. This manuscript was submitted on August 14, 2023; approved on January 8, 2024; published online on April 17, 2024. Discussion period open until September 17, 2024; separate discussions must be submitted for individual papers. This paper is part of the *Journal of Hydrologic Engineering*, © ASCE, ISSN 1084-0699.

⁴Associate Professor, Taylor Geospatial Institute, Saint Louis Univ., St. Louis, MO 63103. ORCID: <https://orcid.org/0000-0003-4375-2096>. Email: vasit.sagan@slu.edu

⁵Associate Professor, Dept. of Civil and Environmental Engineering, Univ. of Iowa, Iowa City, IA 52242. ORCID: <https://orcid.org/0000-0002-0461-1242>. Email: ibrahim-demir@uiowa.edu

⁶Research Engineer, IIHR Hydrosience and Engineering, Univ. of Iowa, Iowa City, IA 52242. Email: marian-muste@uiowa.edu

⁷Hydraulic Engineer, US Army Corps of Engineers, 1616 Capitol Ave., Ste. 9000, Omaha, NE 68102. Email: Paul.M.Boyd@usace.army.mil

⁸Hydrologic and Hydraulic Engineer, US Army Corps of Engineers Headquarters, 441 G St. NW, Washington, DC 20314-1000. Email: Chandra.S.Pathak@usace.army.mil

exceedingly expensive, with dam removal providing the greatest expense in dam decommissioning options (USBR 2006).

Existing reservoir sedimentation models have been unable to analyze the intricate large-scale temporal or spatial patterns of sedimentation due to a lack of available data required for model calibration and validation. The typical data required for model construction include daily to yearly hydrologic records, bathymetric reservoir details, and grain-size distribution of sediment (Ackers 1988; Lajczak 1996; Tarela and Menéndez 1999; Sundborg 1992; Rowan et al. 2001). The most valuable support for reservoir sedimentation model development in recent times has been provided by geographic information system (GIS) tools that enable the addition of land use over large scales to the hydrologic data (Verstraeten et al. 2003; Vörösmarty et al. 2003; Lehner et al. 2011). However, GIS tools are relatively new, hence their historical records are too short to refine sedimentation modeling (Xu et al. 2019). This lack of temporal data in sedimentation modeling diminishes the ability for proper model calibration, which has shown in sediment yield estimated values to deviate considerably from measured sediment yield rates (Trimble 1999).

USACE oversees several dams and reservoirs across the United States, with many being under operation for more than 50 years (Pinson et al. 2016). The aging of these USACE reservoirs puts them at greater risk for complications related to sedimentation. Reservoir capacity surveys focused on US reservoir sedimentation trends indicate that they could deplete by as much as 10%–35% of absolute water storage capacity (Randle et al. 2021). These historical surveys are invaluable tools for identifying past and present regional sedimentation trends, allowing for the evaluation of sediment aggradation and life expectancy of individual reservoirs. These data are also relevant for developing effective reservoir management strategies. Ensuing from the preceding, USACE initiated the Enhancing Reservoir Sedimentation Information for Climate Preparedness and Resilience Program, which includes the Reservoir Sedimentation Information (RSI) system, to assess reservoir aggradation and track dam operation suitability for water resource management. However, because the RSI data set contains less than 1% of US dams, developing methods for estimating reservoir sedimentation at unmonitored sites is needed.

Machine learning as a tool for prediction and anomaly detection has developed rapidly over the past couple of decades. Idrees et al. (2021) developed six machine learning models to predict sediment load inflow into a reservoir located in South Korea using streamflow, water temperature, and reservoir outflow. Aytek and Kişi (2008) proposed a genetic programming model to relate streamflow and suspended sediment values at Tongue River in Montana. Malik et al. (2019) proposed a genetic programming model to predict daily suspended sediments at Godavari River basin, India, using the discharge and the suspended sediment concentration of previous days. Hassan et al. (2022) applied an artificial neural network to estimate the amount of sediment deposited at the Tarbela dam in Pakistan based on yearly rainfall, water inflow, water level at the reservoirs, and storage capacity.

Through several research studies, machine learning has been proven to be successful at predicting streamflow, sediment transport, sediment deposition, and water quality characteristics as well as identifying data anomalies (Xiang and Demir 2020; Azamathulla et al. 2010; Choubin et al. 2018; Peterson et al. 2019, 2020; Xu et al. 2019; Bhadra et al. 2020; Hazarika et al. 2020). Due to the nonlinear behavior of sedimentation processes influenced by various hydraulic flow factors, the use of machine learning has great potential for constructing accurate reservoir capacity loss at unmonitored sites compared with alternative methods (Adnan et al. 2019; Baniya et al. 2019). Machine learning utilizes the process of

iteration and probabilistic pattern detection to determine the relationship between input parameters and a dependent variable (Geron 2022). Prior to utilization of machine learning applications, the sediment yield and sediment load as well as estimated water pollutants were obtained through various process-based models (Ayele et al. 2017; Zounemat-Kermani et al. 2019, 2020).

Machine learning modeling applied to reservoir sedimentation is not, however, infallible as shown by the back-propagation networks used to assess sediment transfer occurring under differing land use and agricultural practices (Abrahart and White 2001). The valuable insights provided by an artificial neural network model trained on 32 years of reservoir sedimentation data for one reservoir (Jothiprakash and Garg 2009) indicate that the availability of long-term data is critical for a trustful modeling outcome. It is, however, obvious that training machine learning requires not only long-term data but also a great variety of reservoirs in order to be more reliable and generalizable.

Although several studies have applied machine learning models to investigate reservoir sedimentation and sediment transport, none, to our knowledge, have implemented these techniques with a data set containing information from multiple reservoirs located within a wide latitude–longitude span. Moreover, these models have primarily used discharge–sediment relationships and have not included other associated watershed variables, such as soil characteristics, land use, or geographical location. Finally, reviewed studies were focused on the estimation of daily suspended sediment loads in streams or reservoir inflow and outflow, and a limited number of studies have been focused on estimating site-specific long-term trends of reservoir capacity losses (Hassan et al. 2022).

The RSI system provides a good baseline resource for training data-driven models that could be utilized for improved reservoir sedimentation estimation modeling through its combination of temporal and spatial data spanning the contiguous United States. The objective of this research was to create a generalized deep learning (DL) method for estimating long-term reservoir sedimentation trends using reservoir design data, historical elevation–capacity records, and supplemental watershed information for reservoirs located across the contiguous United States. To achieve this objective, the following tasks were completed: (1) the RSI data set was analyzed to determine capacity loss between consecutive surveys, (2) supplemental hydrologic data were derived for each reservoir and set of consecutive surveys (e.g., basin area and cumulative precipitation), (3) multiple deep learning algorithms were applied using the composite data set to create models to predict reservoir sedimentation, and (4) model performances were analyzed and compared to identify a recommended model for industry use. This prediction tool will allow the estimation of current conditions of unmonitored reservoirs and forecast future sedimentation rates for reservoirs in the United States.

Composite RSI Data Set Development

RSI information for reservoirs with three or more surveys (184 reservoirs shown in Fig. 1) was combined with supplementary watershed information related to hydrologic and sedimentation processes to form the composite RSI data set utilized in this study. Each record of this data set corresponded to two consecutive surveys conducted by USACE at that particular reservoir, and the capacity loss for each record was the difference in the reservoir's capacity at the maximum pool elevation that was not characterized as a surcharge pool. These computed capacity losses have inherent uncertainty stemming from a wide range of sources that can be broadly grouped by varying factors including (1) field data

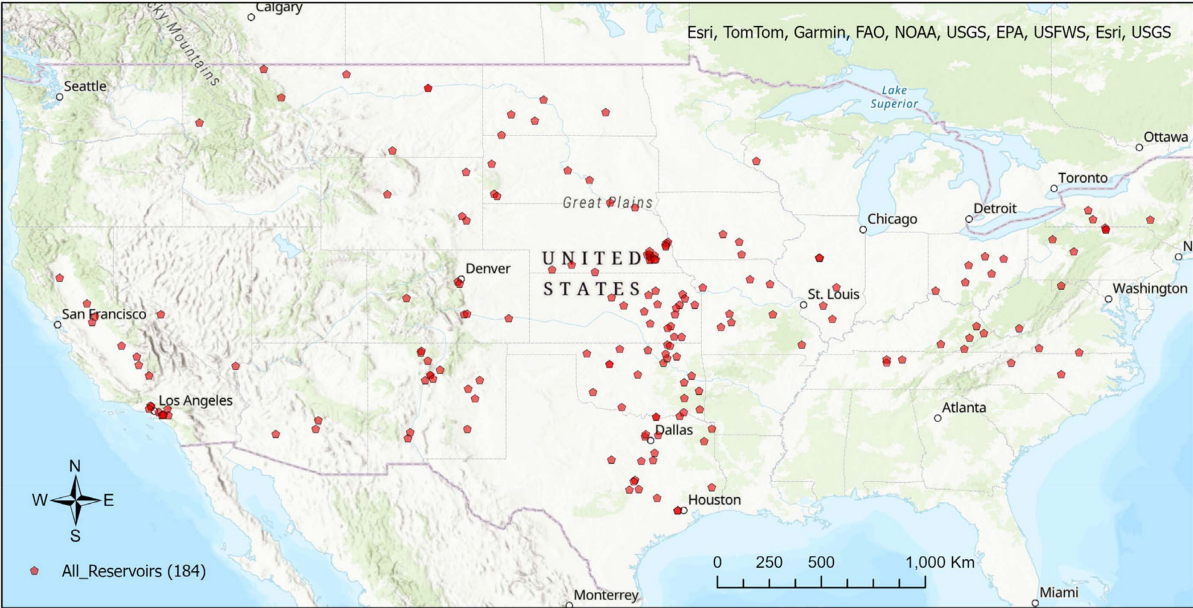


Fig. 1. Map of the 184 reservoirs used in the study. (Base map © Esri, TomTom, Garmin, FAO, NOAA, USGS, EPA, USFWS, Esri, USGS.)

collection instruments, (2) measurement methods, (3) data reduction protocols, (4) operator skills, and (5) measurement environments. A comprehensive analysis of uncertainty was not included in this study due to the limited information available related to these uncertainty sources.

The RSI composite data set incorporated data remotely compiled through publicly available sources to ensure comprehensive watershed characteristics were associated with each recorded reservoir’s capacity loss. Utilization of raster data sets enabled the extraction of relevant hydrologic data, which were identified and applied to their respective basins in the composite database. These public databases provided data related to climatologic, topographic, and erosion processes occurring across the associated watersheds for each record in the composite data set. Fig. 2 shows the features collected per each reservoir record, including the originally provided USACE RSI system data, and the accessed public

database. Reservoir features and basin characteristics in the data set were assumed constant over time for each reservoir. Thirty-two variables compose the composite RSI data set, including numerical variables (27), identifier variables (two), categorical variables (two), and a date variable. Missing data records were replaced with the mean for that specific variable.

The watershed centroid latitude and longitude values for each reservoir were extracted from each basin’s shapefile. The curve number (CN) and the erodibility values were computed for each reservoir as the area-weighted average for its associated basin. The CN is the empirical hydrologic parameter indicative of a catchment’s runoff potential based on soil and land use characteristics (USDA SCS 1986), while the erodibility index is an empirical measure of the inherent resistance of geologic materials (soils and rocks) to erosion. The CN maps were based on national soil and national land cover raster files (Viger and Bock 2014;

Reservoir Sedimentation Information Database	• Latitude • Longitude • Original capacity	• Construction year • Capacity change • Sedimentation rate	• Period between surveys • Time since construction	
Digital Elevation Model	• Basin area • Elevation relief • Max. elevation • Min. elevation	• Mean elevation • Median elevation • Elevation standard deviation	• Mean slope • Hydraulic length • Channel slope	• Initial trap efficiency
National Landcover Database	• Percent forested area			• Curve number
USGS Soil Maps	• Soil type → Erodibility			
Monthly Precipitation Maps	• Mean monthly precipitation • Max. monthly precipitation	• Median monthly precipitation • Normalized max. precipitation	• Cumulative precipitation	
National Inventory of Dams	• Cumulative upstream dam height	• Total upstream normal storage	• Total upstream maximum storage	
EPA Classification	• U.S. ecoregions			
IECC Classification	• U.S. climate zones			

Fig. 2. Data sources and derived variables (numerical and categorical) of the composite RSI data set.

USGS 2017). Utilizing USGS characteristics for soil hydrologic groupings and land use categorization, the CN values were defined based on guidelines found in the Revised Universal Soil Loss Equation (RUSLE) for each soil type (Renard 1997). The average erodibility indexes for sand (0.125), loam (0.325), and clay (0.1) were used to create erodibility maps for each reservoir's basin based on the national soils map (Viger and Bock 2014). Additionally, the NLCD was used to compute the percent of forested area within a reservoir's basin, with deciduous, evergreen, and mixed forest types consolidated into one category for this study.

The Google Earth Engine facilitated the extraction and computation of variables from US digital elevation models (DEMs) and monthly precipitation maps (USGS 2017; Gorelick et al. 2017). For this analysis, a 1/3 arcsecond DEM was utilized for calculating features reliant on topographic information for the 184 reservoir basins in the composite data set. These features include hydraulic length, basin elevation, average slope, area, and relief, which was defined as the difference between the maximum and minimum elevation. Based on these calculations, the channel slope was estimated as the relationship of the relief divided by the hydraulic length. A reservoir's initial trap efficiency (E) was calculated as a reservoir's initial capacity in cubic meters (C), and a reservoir's drainage area in square kilometers (A) as shown in Eq. (1) (Brown 1943)

$$E = 1 - \frac{1}{1 + (2.1 \times 10^{-4})C/A} \quad (1)$$

Further, precipitation data for each reservoir were found by analyzing 30 arcsecond monthly precipitation raster files (Daly et al. 2015) that aligned with the database's time periods per each set of consecutive surveys. Additionally, cumulative, maximum, mean, and median monthly precipitations for each record were calculated. Further, the computation of normalized maximum precipitation equaled the maximum precipitation divided by the mean monthly precipitation.

Because many dams were built upstream of RSI reservoirs, a batch analysis was employed to include upstream dam heights, as well as the maximum and normal storage of each reservoir in the RSI composite data set. This computation was conducted in two steps: (1) Utilize the National Inventory of Dams (NID) data set, composed of more than 90,000 US dams, to create an annual time series of cumulative upstream dam height, and normal and maximum storage for each RSI reservoir; and (2) time average the upstream dam's variables for the period of time comprising two subsequent surveys for each RSI data set record.

Data Set Preprocessing

Due to natural processes, sustained or increases in reservoir capacity are not possible, unless dredging or free-flow sediment flushing has been employed (Wang and Hu 2009). Thus, a reservoir's capacity will decrease over time. With this knowledge, the RSI composite data set records containing identical capacities or an increased trend in capacity between a set of consecutive surveys were removed. Additionally, sets of consecutive surveys containing identical survey data or dates were filtered out. Applying the preprocessing methods resulted in a composite data set containing 467 records (sets of consecutive surveys) from 174 reservoirs located across the United States. Each record represents a unique pair of subsequent surveys performed at a specific reservoir.

A log transformation (Brakstad 1992; Emmerson et al. 1997) was applied to the numerical variables of the RSI composite data

set to remove the impact of the difference in orders of magnitude. The following provides the equation for the log transformation:

$$x_{li} = \text{sgn}[\ln(|x_i| + 1)] \quad (2)$$

where x_i = original data value; x_{li} = log-transformed value; i = number of observations; and the sgn function multiplies the value by either a value of 1 if x_i is a positive value or a value of -1 if x_i is a negative value. Additionally, a minimum-maximum (min-max) normalization (Goyal et al. 2014; Patro and Sahu 2015) of the numerical variables was conducted using the following equation:

$$x_{lmi} = 0.7 \left(\frac{x_{li} - x_{l_min}}{x_{l_max} - x_{l_min}} \right) + 0.15 \quad (3)$$

where x_{lmi} = log-transformed and min-max normalized value; x_i = original data value; x_{li} = log-transformed value; x_{l_min} = minimum value of the x_l data set; and x_{l_max} = maximum value of the x_l data set. This results in a linear scaling with values ranging from 0.15 to 0.85. The min-max normalization of data fits the data in a predefined range keeping the relationships from the original data unchanged (Patro and Sahu 2015).

Depending on the performance of models, standard scaling was applied in lieu of the min-max normalization. This normalization method minimizes the number of parameters that appear constant across the data set, which can affect model performance. Standard scaling centers the data set values around the mean with a unit of standard deviation (Cao et al. 2016). The following equation details the standard scaling calculations:

$$x_{lsi} = \frac{x_{li} - \mu}{\sigma} \quad (4)$$

where x_{lsi} = log-transformed scaled value; μ = mean of the x_{li} data set; and σ = standard deviation of the x_{li} data set.

Methods

The data set compiled for the RSI reservoir sites consisted of variables relevant to sedimentation and hydrologic processes. Transformation and scaling of the data set were performed to diminish bias and skew of the variables' distribution. A feature importance analysis was conducted to analyze the sensitivity of variables detrimental to model performance, which resulted in the creation of a data set with decreased variable size. The original and the feature-importance-derived data sets were used to develop and evaluate capacity loss prediction models. Both sets of data were examined in each iteration of the statistical or machine learning method. For all models analyzed, a 70/30 split of the data set was applied for the training and testing of the models, respectively.

The first statistical model used was the ordinary least squares (OLS) multilinear regression model. The second analysis consisted of four supervised machine learning regression models: support vector machine (SVM), random forest (RFR), decision tree (DTR), and partial least squares (PLS). The third analysis used deep neural network (DNN) models. In the DNN model survey, four base DNN architectures were analyzed.

A data anomaly detection was performed to reduce erroneous data in the composite data set. This included anomaly removals utilizing autonomous anomaly detection (AAD) (Angelov et al. 2016; Gu and Angelov 2017), which flagged 18 records corresponding to 15 reservoirs, and the Kolmogorov-Smirnov and Efron (KSE) outlier detection method (Jirachan and Priomsopa 2015), which flagged 15 records corresponding to 10 reservoirs. Removal

of anomalous data from data used in model development varied by model based on performance.

Last, seven metrics were used to compare all created models. The following performance parameters were quantified for evaluation and goodness-of-fit analysis of the statistical models: coefficient of determination (R^2), mean absolute percent error (MAPE), root-mean-square error (RMSE), and relative root-mean-square error (RRMSE). The remaining three parameters included the percent bias (PBias), the ratio of root-mean-squared error to standard deviation of measured data (RSR), and the Pearson correlation coefficient (r) to help analyze the models' overall accuracy outside the limitations of correlation-based measures (Legates and McCabe 1999). Respectively, these three metrics were used to quantify each model's overestimation or underestimation and normalization to error index evaluation in model performance (Moriasi et al. 2007), and uncover the degree of linear association between calibrated and observed values of the model (Taylor 1990; Adler and Parmryd 2010). Collectively, watershed model performance metrics can be considered satisfactory if $R^2 > 0.5$, PBias $< \pm 55\%$, RSR < 0.7 , and $r > 0.7$ (Moriasi et al. 2007; Ayele et al. 2017).

Feature Correlation and Recursive Feature Elimination

A Spearman's rank correlation calculation was performed to measure the monotonic relationship across predictor variables. Ranging from -1 to 1 , the Spearman's calculated coefficient gauges whether two features are correlated, with -1 being negatively correlated and 1 being positively correlated (Bon-Gang 2018). Determining these relationships between the predictor variables was necessary to investigate potential collinearity shared across the composite data set, and if removal of features could improve ensuing model performance. The general criterion for modeling a regression analysis is a minimum of 10 to 20 samples per predictor variable (Austin and Steyerberg 2015).

To observe if reducing the number of predictor variables in the composite data set improved results, a recursive feature elimination (RFE) algorithm was performed from which an alternative data set was developed. Utilizing optimized random forest model parameters, the RFE was used to establish the optimal amount of predictor variables for this new RFE-determined data set. The RFE algorithm assigns weights to features based on model performance. The significance of this algorithm is its allowance to choose the number of features desired in the reduced data set, and its theoretical improvement in statistical modeling through its removal of collinear features. The presence of numerous collinear features can lead to

overfitting when analyzing the prediction of dependent variables through machine learning models (Harrell 2001).

Ordinary Least Squares Multilinear Regression

The OLS multilinear regression model is used for relational analysis between one or more variables. The method corresponds to the minimization of the sum of the square error difference between the observed and predicted values of the target variable because it fits an assumed linear relationship between the explanatory variables (Zdaniuk 2014). The OLS regression formula to compute capacity loss, y_i , is

$$y_i = \beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \cdots + \beta_p x_{ip} + \epsilon \quad (5)$$

where i = number of observations; y_i = dependent variable; x_{ij} = explanatory variables; β_0 = y-intercept or constant term of the equation; β_p = slope coefficients for each explanatory variable; and ϵ = residuals of the model (Alexopoulos 2010). Standard scaling was used to evaluate the magnitude of influence each predictor variable had on the target variable. Imperial system units were used for the OLS analysis and, due to regression model complexity, all International System of Units (SI) values must be converted to imperial units for application. When using the OLS method with metric units, the SI value of each input parameter should be multiplied by the corresponding metric unit conversion factor listed in Table 1 prior to the log transformation in Eq. (2). Similarly, the capacity loss term, y_i , computed from Eq. (5) is in acre-feet and needs to be multiplied by 1,233 for conversion to cubic meters.

Supervised Machine Learning

Supervised machine learning is the application of algorithms capable of producing generalities in patterns via the use of externally supplied data to predict future patterns and instances (Singh et al. 2016). Several types of supervised machine learning algorithms exist, but for this analysis SVM (Noble 2006), RFR (Breiman 2001), DTR (Bashar et al. 2019), and PLS (Manikanta and Mamatha Jadav 2015) regression algorithms were utilized. Each algorithm has advantages and disadvantages when applied to a unique data set; thus, implementation of these four enabled comprehensive analysis of supervised learners on the composite and RFE data sets. Additionally, each supervised learner used a pipeline of several intermediary steps that chained a sequence of estimators for optimization and cross-validation of model performance. These steps included a principal component analysis and

Table 1. RFE ranked data set variables

Index	Variable	Imperial units (SI)	Metric unit conversion factor	Calibrated standard scaled data, OLS coefficients	Calibrated unscaled data, OLS coefficients
1	Basin area	mi ² (km ²)	0.386	1.42	0.553
2	Initial capacity	acre-ft (m ³)	8.11×10^{-4}	1.03	0.476
3	Cumulative precipitation	in. (mm)	3.93×10^{-2}	0.323	0.383
4	Hydraulic length	ft (m)	3.28	-0.259	-0.181
5	Max monthly precipitation	in. (mm)	3.93×10^{-2}	0.234	0.561
6	Curve number	N/A	—	0.144	1.63
7	Total upstream dam height	ft (m)	3.28	-0.119	-0.0494
8	Total upstream normal storage	acre-ft (m ³)	8.11×10^{-4}	0.100	0.0250
9	Basin relief	ft (m)	3.28	-0.0369	-0.0267
10	Channel slope	ft/ft (m/m)	1.00	0.0226	1.91
11	Average basin latitude	Degrees	—	0.0197	0.192
12	Mean monthly precipitation	in./month (mm/month)	3.93×10^{-2}	0.0158	0.0589

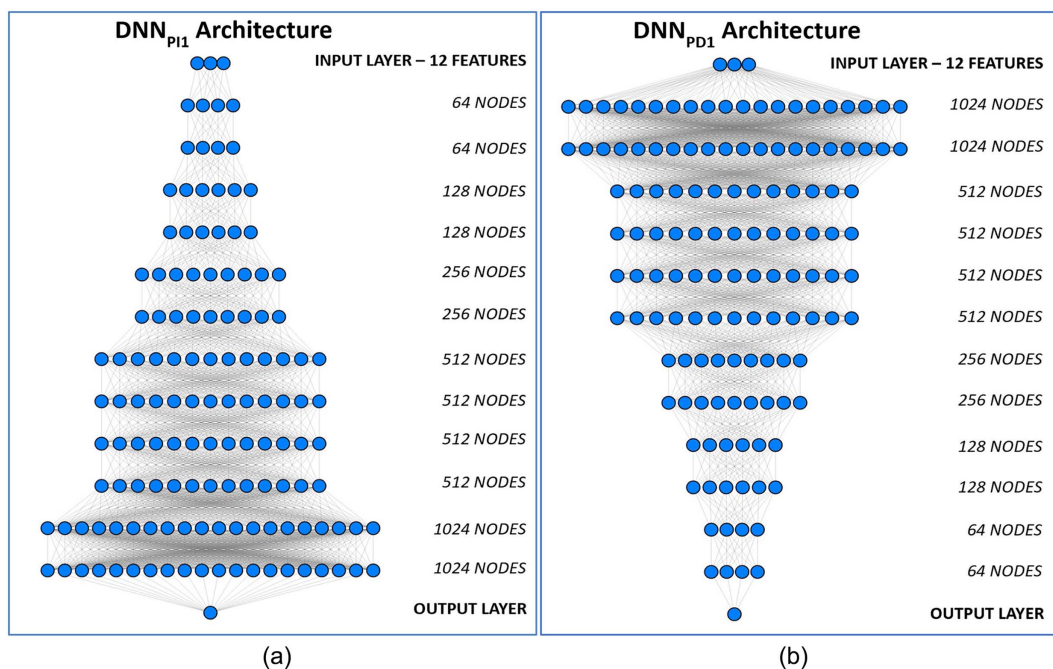


Fig. 3. Hidden-layer architectures of (a) DNN_{P11} ; and (b) DNN_{PD1} with the respective nodes present in each of their layers.

standard scaling. These optimizing components were refined by automated selection of each model's hyperparameters resulting in the highest performing variation of the supervised model.

Deep Neural Network

A DNN is an organized collection of neurons sequenced into multiple layers for determining modeled predictions. The neurons receive input from the initial data set if they reside in the first layer of the DNN, or from input from activated neurons from previous layers if residing in a subsequent layer. The activations of the neurons occur based on a calculation of the weighted sums from that input followed by a nonlinear activation (Montavon et al. 2018). In the case of this analysis, the rectified linear (ReLU) activation function was used. All nodes that consist of this activation function are considered rectified linear activation units (ReLU), whose development was a milestone in the evolution of deep learning (Goodfellow et al. 2016). DL studies have gained significant momentum with the availability of computational resources, benchmark data sets (Demiray et al. 2021; Sit et al. 2021b), and the popularity of DL algorithms in many data analysis tasks in water resources and hydrology including streamflow forecasting (Sit et al. 2021a), culvert sedimentation (Xu et al. 2019), data augmentation (Demiray et al. 2021), and image synthesis (Gautam et al. 2022).

For this analysis, four DNN architectures were utilized, and the models were optimized to minimize the mean absolute error (MAE). The basis of the first DNN architecture was used in the research of Maimaitijiang et al. (2020), which contained a GIS and remotely sensed data set. Named DNN-F1, it incorporated a DNN node structure that continually increased in complexity per each layer. The minimum number of nodes residing in the initial layer was 64, and the maximum number of nodes retained in the final layering was 1,024. For the purposes of this study, this first progressively increasing DNN (termed DNN_{P11}) will be the base DNN used to compare further DNN architectures. The second DNN architecture was aimed at analyzing if information bottlenecks could improve the initial DNN. The bottlenecks method aims to balance improved accuracy through decreasing complexity

(Tishby et al. 2000; Hecht and Tishby 2005). This version of the DNN reverses the initial architecture to become a progressively decreasing DNN (termed DNN_{PD1}), which results in its initial layer containing a node network of 1,024, and its final layer containing a node network of 64. Schematics of the DNN_{P11} and DNN_{PD1} architectures are shown in Fig. 3.

Two simplified DNN structures were also evaluated to determine their performance compared with the complex DNN_{P11} and DNN_{PD1} structures: a second progressively increasing DNN (termed DNN_{P12}) and a second progressively decreasing DNN (termed DNN_{PD2}). Detailed node architectures for these simpler DNNs are shown in Fig. 4. The DNN_{P12} and DNN_{PD2} structures have half the number of layers as the DNN_{P11} and DNN_{PD1} , and fewer nodes associated with each of their layers. The DNN_{P12} structure contains an initial neural structure that starts with eight nodes and increases to 32 in its final layer. The DNN_{PD2} structure is the reversed iteration of DNN_{P12} .

Results and Discussion

Feature Importance Analysis

To identify collinearity or monotonic relationships between features, a Spearman's rank coefficient matrix analysis was performed. Values closest to 1 or -1 were respectively deemed highly positively or negatively correlated. The DEM parameters showed a significant positive correlation, as well as basin relief with values ranging from 0.57 to 0.92. Alternatively, the DEM parameters appear negatively correlated with the monthly precipitation parameters with values of -0.54 to -0.66. This analysis signifies that the compiled features in the data set contain redundancies.

Recursive Feature Elimination

The RFE algorithm was applied to reduce potentially redundant features and further optimize the performance of the predictive models. The composite data set consisted of 467 samples with

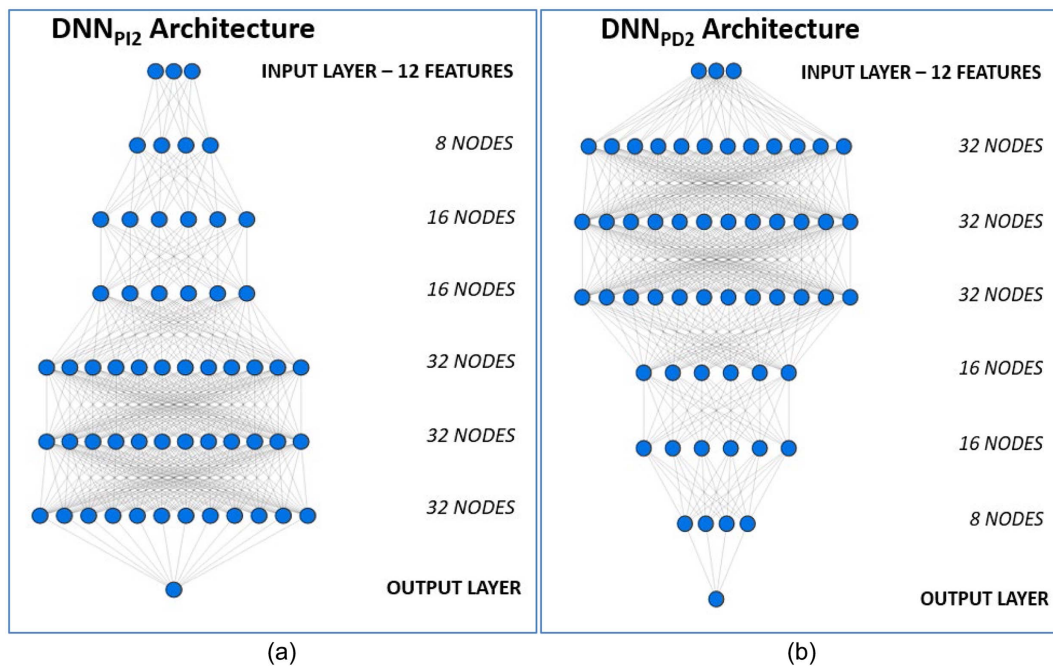


Fig. 4. Hidden-layer architectures of (a) DNN_{p12}; and (b) DNN_{pd2} with the respective nodes present in each of their layers.

27 predictor variables. Thus, the data set had a ratio of approximately 17 samples per predictor variable in the data set. When reducing composite data set features, the RFE conducts its model accuracy performance based on R^2 values, with 1.00 being the highest accuracy score possible. The RFE results showed that 12 predictor variables retained an R^2 value of between 0.78 and 0.80. With less than 12 variables, the accuracy scored less than or equal to 0.77. Thus, the 12 predictor variables listed in Table 1 were optimal in minimizing the composite data set to a sample-to-feature ratio of approximately 38. This new RFE data set was used in subsequent models and the results were compared with the entire composite data set. Standard scaling was used to help further analyze the magnitude of influence each predictor variable had on the target variable in the OLS equation.

The inclusion of basin relief, hydraulic length, and the channel slope features in the RFE data set may be seen as still maintaining excessive collinear features. However, due to the logarithmic transformation and normalizations performed on the data, the feature of channel slope, which is derived from hydraulic length and basin relief, is mathematically unique in terms of providing a predictive value in the model's equation.

Three features are indicators of drainage basin size: basin area, hydraulic length, and basin relief. The basin area model coefficient of 0.548 indicates that area is the dominant feature related to basin size. Based on the Spearman correlation analysis, both basin length and relief are positively correlated with capacity loss. However, both the length and relief coefficients are inversely related to the predictive variable (i.e., capacity loss), suggesting their model contribution is an adjustment on the basin area influence.

All models developed using the RFE data set resulted in improved performance compared with models from the entire composite data set. Due to the large number of models generated, only the RFE results are reported.

OLS Regression Model

The log-transformed RFE data set with standard scaling and no anomalies removed was found to produce the best OLS model

performance. The observed versus predicted training and testing results for the OLS model are shown in Fig. 5. Following the training and testing analysis, the OLS model was calibrated using the full data set to provide the overall best-fit equation. Fig. 6 shows the observed versus predicted values for the calibrated OLS model that had an R^2 value of 0.40 and a MAPE of 195%. Eq. (6) provides the OLS prediction equation based on the coefficient values and constant terms derived from the calibrated OLS model results

$$y_{OLS_i} = -9.71 + 0.548x_{i1}U_1 + 0.476x_{i2}U_2 + 0.383x_{i3}U_3 - 0.169x_{i4}U_4 + 0.561x_{i5}U_5 + 1.59x_{i6}U_6 - 0.0460x_{i7}U_7 + 0.0250x_{i8}U_8 - 0.0249x_{i9}U_9 + 1.87x_{i10}U_{10} + 0.188x_{i11}U_{11} + 0.0588x_{i12}U_{12} \quad (6)$$

where y_{OLS_i} = log-transformed predicted capacity; x_{ip} = log-transformed predictor variables; U_p = metric unit conversion factor; and the numeric subscript p on the x_i and U terms denotes the variable index (Table 1). To obtain the predicted capacity loss value, the model-predicted value (y_{OLS}) needs to be untransformed using Eq. (7)

$$y_{OLS} = (e^{y_{OLS_i}} - 1) \times 1,233 \quad (7)$$

Supervised Machine Learning

The best performing supervised machine learning model was identified based on the satisfactory statistical metrics defined by Moriasi et al. (2007). Nearly all the supervised machine learning models had optimal performance when using the log-transformed min-max normalized RFE data set with the KSE anomalies removed. The supervised machine learning results presented in this paper were all developed using this data set. A comparison between the supervised machine learning methods showed that RFR had the most accuracy, in terms of predictive performance, when trained and tested on the respective data. With a training set R^2 of 0.61 and a testing set R^2 of 0.57, the model shows precision in model fitness when comparing the predicted versus observed values of

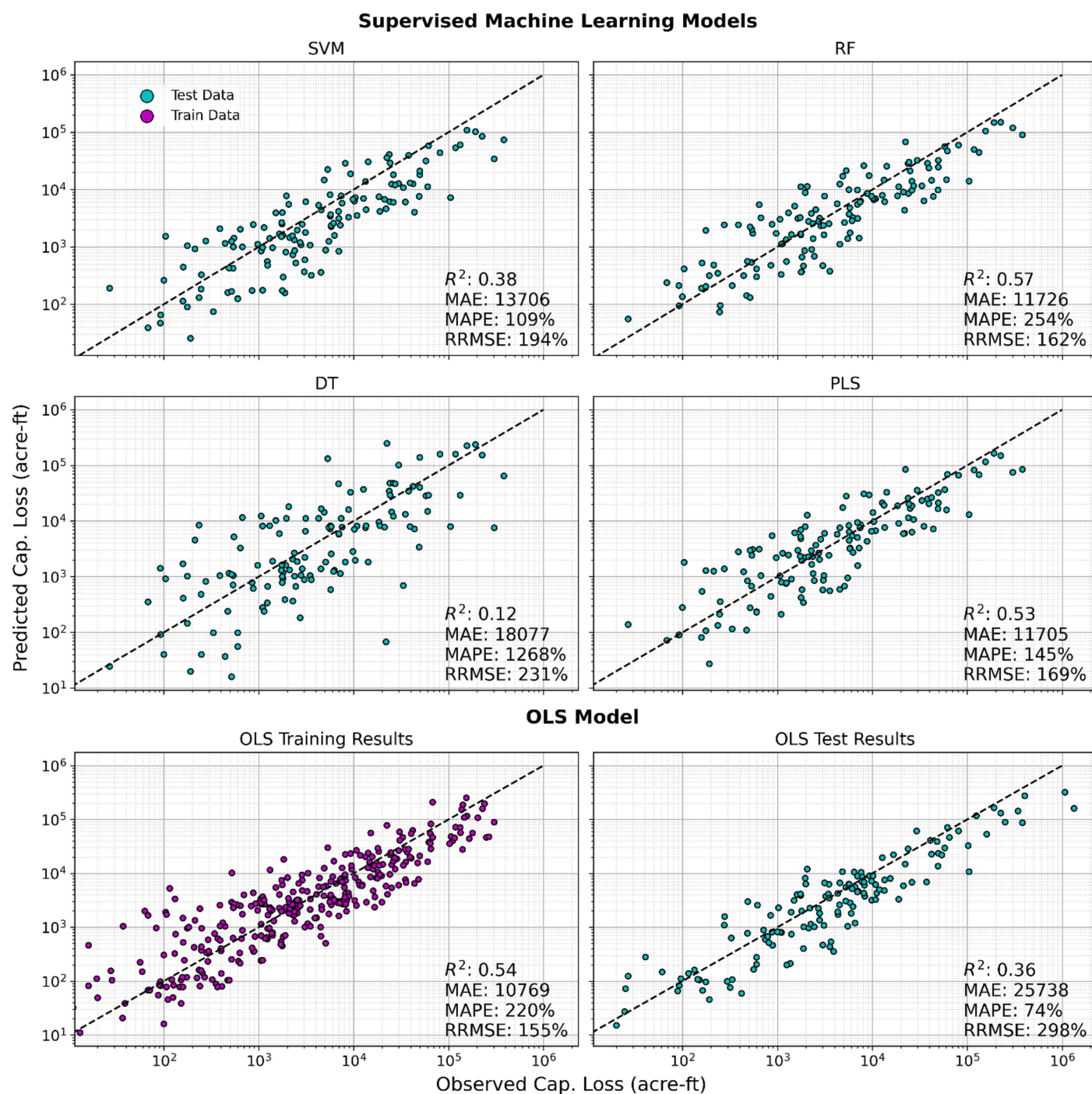


Fig. 5. Comparison of supervised machine learning and OLS models.

capacity loss. Figs. 7 and 8 show the performance metrics of the training and testing results, respectively, for the RFE model. Notably, there is a significant increase in MAPE on the testing data set's forecasting accuracy. This signifies that the model training results are overestimating the model's performance, regardless of the relatively high R^2 value present on the testing data set.

DNN Analysis

All the DNN models had optimal performance when using the log-transformed min-max normalized RFE data set with the KSE anomalies removed. The DNN results presented for this study were

developed using this data set. The complex DNNs had significantly better accuracy based on the MAPE and R^2 values. The DNN_{PII} was identified as the best DNN model variation based on maximizing the R^2 and minimizing the RRMSE. Training and testing results for this DNN model are shown in Figs. 7 and 8, respectively. The DNN_{PII} had training and testing R^2 values of 0.83 and 0.70, respectively. This makes the DNN_{PII} the best fitting model in terms of performance. The RRMSE values of the DNN_{PII} were the lowest RRMSE values across all analyzed machine learning models. However, the MAPE and RRMSE values showed a relatively large percentage increase between training and testing, meaning there may be underlying forecasting inaccuracies.

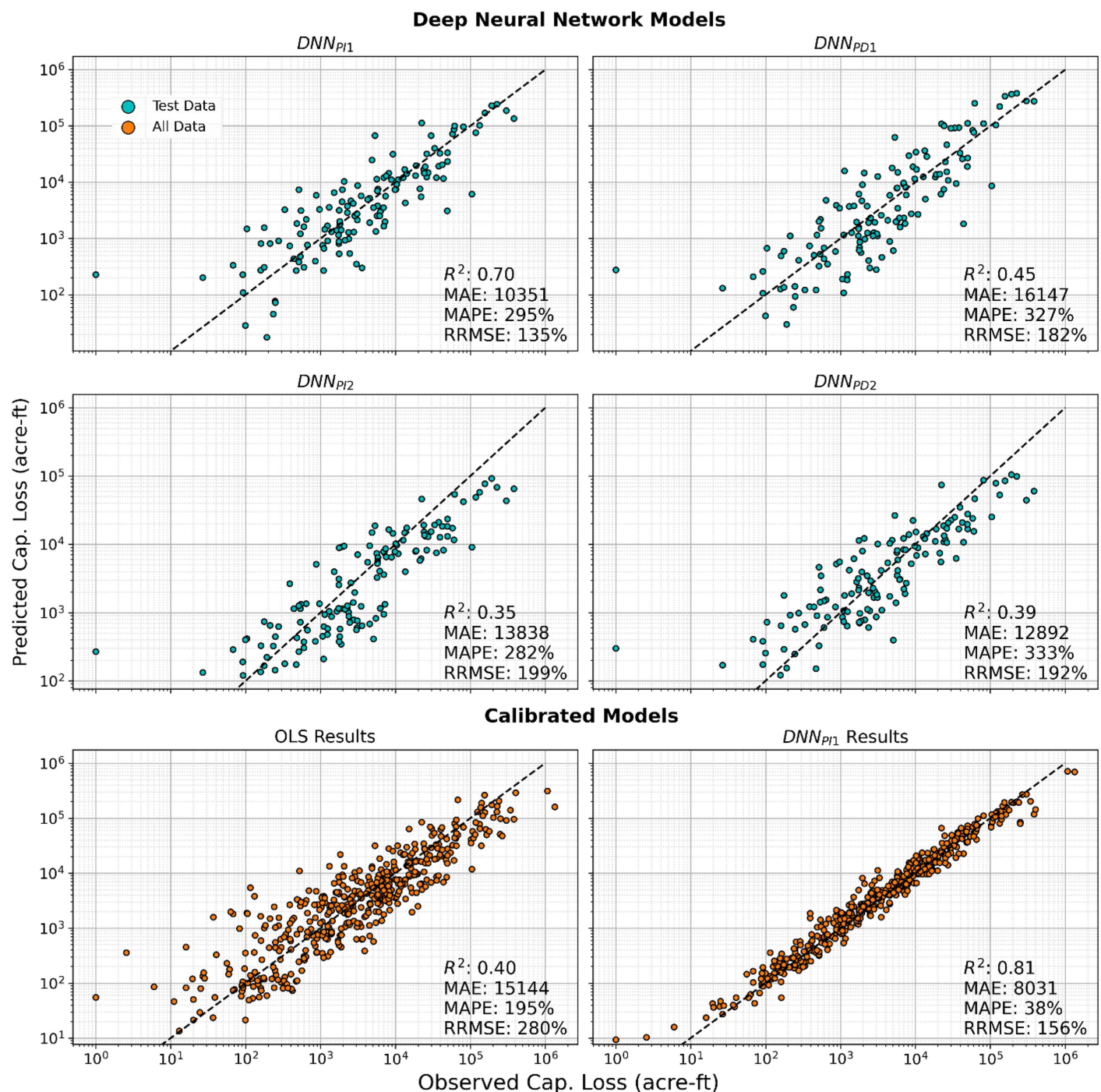


Fig. 6. Comparison of DNN models and calibrated DNN and OLS models.

Comparison of Models

Comparisons of the supervised machine learning, DNN, and OLS models are shown in Figs. 5 and 6; model summary statistics for the untransformed model data are provided in Figs. 7 and 8. All models were developed using the transformed data, but the prediction variable of interest is the capacity loss (i.e., not the log-transformed capacity loss). Thus, untransformed statistics were used to assess model performance and their values are reported on all observed versus predicted plots. Further, results shown are for the RFE data set (feature variables listed in Table 1) because the RFE data set performed better than the original composite data set for all models analyzed. Except for RFR, the supervised machine learning

methods resulted in abnormal predictive performance. Notably, DTR showed extreme overfitting characteristics, where the model appeared to learn perfectly in terms of its performance metrics on the training data set, but failed to accurately replicate this performance on the testing data set. However, RFR, and the more complex DNNs, showed promising results in terms of learning and predicting capacity loss. Overall, the best tested model performance, based on R^2 and RRMSE, was the DNN_{PI1} with an untransformed R^2 value of 0.70 and an untransformed RRMSE of 135%.

The RRMSE values measured across all models as they relate to the OLS RRMSE value are shown in Fig. 9. The OLS method of prediction compared respectably when set side by side with more

MODEL	Untransformed Training Statistics							
	MAE	MAPE	RRMSE	RMSE	R ²	PBIAS(%)	RSR	r
OLS	1.3×10^7	220	155	3.6×10^7	0.54	29.45	0.67	0.75
SVM	2.5×10^7	165	323	1.2×10^8	0.24	59.57	0.87	0.68
RFR	1.4×10^7	40	232	8.2×10^7	0.61	34.35	0.62	0.92
DTR	0	0	0	0	1.00	0	0	1.00
PLS	2.1×10^7	211	294	1.1×10^8	0.38	41.87	0.79	0.71
DNN _{PI1}	1.4×10^7	87	155	5.5×10^7	0.83	14.77	0.42	0.94
DNN _{PD1}	2.3×10^7	106	198	7.0×10^7	0.72	-32.77	0.53	0.88
DNN _{PI2}	2.5×10^7	187	339	1.2×10^8	0.17	64.41	0.91	0.64
DNN _{PD2}	2.1×10^7	189	287	1.0×10^8	0.41	52.15	0.77	0.84
Calib. OLS	1.9×10^7	195	280	8.9×10^7	0.40	39.03	0.78	0.71
Calib. DNN _{PI1}	8.0×10^3	38	156	4.1×10^4	0.81	16.57	0.44	0.95

Fig. 7. Untransformed training metrics summary for capacity loss models; highlighted cells indicate satisfactory metrics as determined by criteria from Moriasi et al. (2007), the Calib. OLS and Calib. DNN_{PI1} models were calibrated with the entire data set.

MODEL	Untransformed Testing Statistics							
	MAE	MAPE	RRMSE	RMSE	R ²	PBIAS(%)	RSR	r
OLS	3.2×10^7	74	298	1.5×10^8	0.36	57.79	0.80	0.80
SVM	1.7×10^7	109	194	5.1×10^7	0.38	54.20	0.79	0.78
RFR	1.4×10^7	254	162	4.2×10^7	0.57	40.47	0.66	0.85
DTR	2.2×10^7	1268	231	6.0×10^7	0.12	9.15	0.94	0.50
PLS	1.4×10^7	145	169	4.4×10^7	0.53	35.87	0.69	0.79
DNN _{PI1}	1.3×10^7	295	135	3.5×10^7	0.70	12.49	0.55	0.84
DNN _{PD1}	2.0×10^7	327	182	4.8×10^7	0.45	-37.95	0.74	0.85
DNN _{PI2}	1.7×10^7	282	199	5.2×10^7	0.35	56.84	0.81	0.82
DNN _{PD2}	1.6×10^7	333	192	5.0×10^7	0.39	46.36	0.78	0.76

Fig. 8. Untransformed testing metrics summary for capacity loss models; highlighted cells indicate satisfactory metrics as determined by criteria from Moriasi et al. (2007).

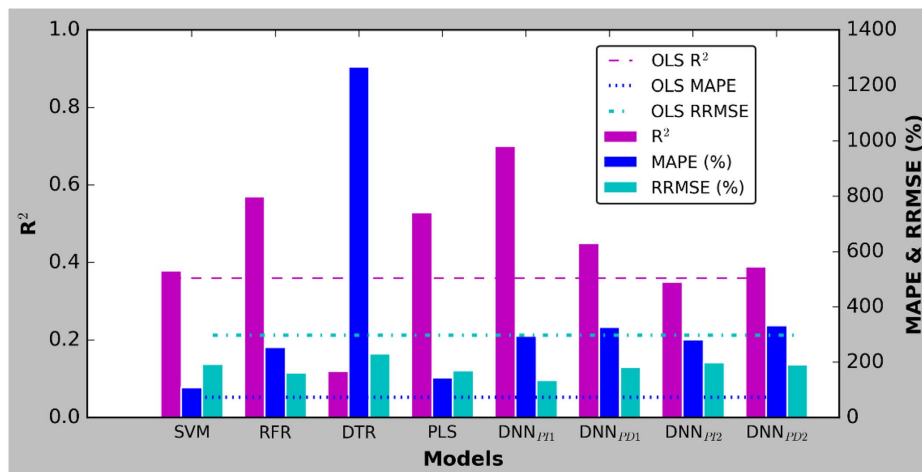


Fig. 9. Comparison of R², MAPE, and RRMSE values across all models, related to the respective OLS method values.

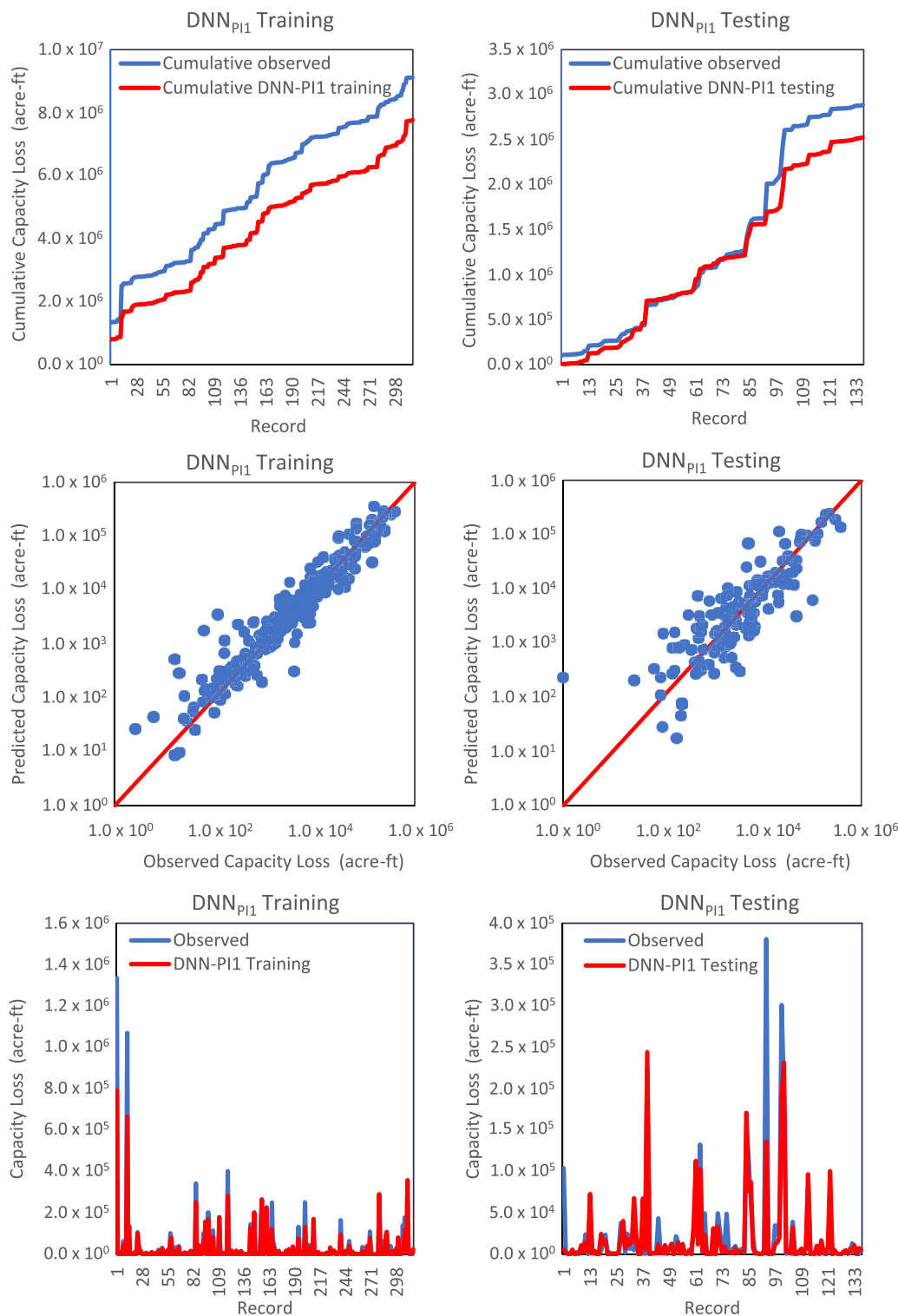


Fig. 10. Cumulative capacity loss, observed versus simulated capacity loss, and capacity loss data series corresponding to the untransformed metrics for the DNN_{PI1} machine learning model.

computationally complex machine learning models in terms of R^2 and MAPE. However, the more complex models did result in considerably lower RRMSE values compared with the OLS method.

As shown in Fig. 7 (training) and Fig. 8 (testing), the DNN_{PI1} model was further proved as the best model because it exhibited satisfactory performance for all training and testing metrics (i.e., R^2 , PBias, RSR, and r). Fig. 10 shows the cumulative capacity loss, observed versus simulated capacity loss, and capacity loss

data series for the DNN_{PI1} model. Across the entire testing data set records, the DNN_{PI1} model estimated 12.5% less than the observed cumulative measured quantities. However, despite this error margin, the model successfully learned from the training data set to produce satisfactory performance metrics based on criteria defined by Moriasi et al. (2007).

Consequently, the model recommended for capacity loss prediction is a calibrated DNN_{PI1} model. The calibrated DNN_{PI1} was

established through training the original best-performing DNN_{PI1} model on the entire RFE data set. This was conducted to overcome potential inaccuracies associated with the limited records available, which is the case with the current RSI data set. For this calibrated model, the R^2 increased to 0.81 and the MAPE value decreased to 38%, as shown in Fig. 7. This shows significant improvement in terms of forecasting accuracy compared with all models. Fig. 6 illustrates the observed versus predicted capacity loss values for the calibrated DNN_{PI1}. Thus, the model successfully learned on the training data set, producing satisfactory performance metrics. However, high-accuracy determinations for larger amounts of capacity loss still appear limited.

Conclusions

A composite data set was developed that included capacity loss data obtained from RSI system records and 29 supplemental parameters derived from publicly available databases. The composite data set included 184 reservoirs, 799 surveys, and 615 sets of consecutive surveys for evaluating capacity loss. The study demonstrated that prediction models containing supplemental data inputs estimate reservoir capacity loss (acre-ft) with satisfactory R^2 , PBias, RSR, and r values as defined in Moriasi et al. (2007). Of the nine predictive models, the progressively increasing deep neural network (DNN_{PI1}) had the best predictive performance with model training and testing R^2 values of 0.83 and 0.70, respectively; and training and testing MAPE of 87% and 295%, respectively. The MAPE performance measured the average magnitude of error produced by the model (87%) on the test data set as larger than the training data set (295%). The DNN_{PI1} model was recalibrated over the entire data set with resulting R^2 and MAPE values of 0.81 and 48%, respectively. Accordingly, the DNN_{PI1} is the most promising model for estimating reservoir capacity losses using watershed and historical precipitation data, which enables the identification of vulnerable reservoirs in the United States. Further, the DNN_{PI1} model can be used to forecast reservoir sedimentation rates under possible future climate scenarios, which allows for the development of proactive management plans.

Data Availability Statement

The fully calibrated DNN_{PI1} model code is available on GitHub (<https://github.com/uihilab/RSI-DNNModel>). The RSI data set used to develop the models is not currently publicly available. All other data used in the study were derived from publicly available data sets.

Acknowledgments

This research was supported by the US National Science Foundation (Award #1948940) and the WATER Institute at Saint Louis University.

References

Abrahart, R. J., and S. M. White. 2001. "Modelling sediment transfer in Malawi: Comparing backpropagation neural network solutions against a multiple linear regression benchmark using small datasets." *Phys. Chem. Earth Part B* 26 (1): 19–24. [https://doi.org/10.1016/S1464-1909\(01\)85008-5](https://doi.org/10.1016/S1464-1909(01)85008-5).

Ackers, P. 1988. "Alluvial channel hydraulics." *J. Hydrol.* 100 (1–3): 177–204. [https://doi.org/10.1016/0022-1694\(88\)90185-0](https://doi.org/10.1016/0022-1694(88)90185-0).

Adler, J., and I. Parmryd. 2010. "Quantifying colocalization by correlation: The Pearson correlation coefficient is superior to the Mander's overlap coefficient." *Cytometry Part A* 77 (8): 733–742. <https://doi.org/10.1002/cyto.a.20896>.

Adnan, M. S., A. Dewan, K. E. Zannat, and A. M. Abdullah. 2019. "The use of watershed geomorphic data in flash flood susceptibility zoning: A case study of the Karnaphuli and Sangu river basins of Bangladesh." *Nat. Hazards* 99 (1): 425–448. <https://doi.org/10.1007/s11069-019-03749-3>.

Alexopoulos, E. C. 2010. "Introduction to multivariate regression analysis." *Hippokratia* 14 (1): 23–28.

Angelov, P., X. Gu, D. Kangin, and J. Principe. 2016. "Empirical data analysis: A new tool for data analytics." In *Proc., 2016 IEEE Int. Conf. on Systems, Man, and Cybernetics (SMC)*, 52–59. New York: IEEE.

Austin, P. C., and E. C. Steyerberg. 2015. "The number of subjects per variable required in linear regression analyses." *J. Clin. Epidemiol.* 68 (6): 627–636. <https://doi.org/10.1016/j.jclinepi.2014.12.014>.

Ayele, G. T., E. Z. Teshale, B. Yu, I. D. Rutherford, and J. Jeong. 2017. "Streamflow and sediment yield prediction for watershed prioritization in the Upper Blue Nile River Basin." *Water* 9 (10): 782. <https://doi.org/10.3390/w9100782>.

Aytek, A., and Ö. Kisi. 2008. "A genetic programming approach to suspended sediment modeling." *J. Hydrol.* 351 (3–4): 288–298. <https://doi.org/10.1016/j.jhydrol.2007.12.005>.

Azamathulla, H. M., A. A. Ghani, C. K. Chang, Z. A. Hasan, and N. A. Zakaria. 2010. "Machine learning approach to predict sediment load—A case study." *CLEAN—Soil Air Water* 38 (10): 969–976. <https://doi.org/10.1002/clen.201000068>.

Baniya, B., Q. Tang, X. Xu, G. G. Haile, and G. Chhipi-Shrestha. 2019. "Spatial and temporal variation of drought based on satellite derived vegetation condition index in Nepal from 1982–2015." *Sensors* 19 (2): 430. <https://doi.org/10.3390/s19020430>.

Bashar, S. S., M. S. Miah, A. Zaidul Karim, and M. Al Mahmud. 2019. "Extraction of heart rate from PPG signal: A machine learning approach using decision tree regression algorithm." In *Proc., 4th Int. Conf. on Electrical Information Communication Technology (EICT)*, 1–6. New York: IEEE.

Bhadra, S., V. Sagan, M. Maimaitijiang, M. Maimaitiyiming, M. Newcomb, N. Shakoor, and T. C. Mockler. 2020. "Quantifying leaf chlorophyll concentration of sorghum from hyperspectral data using derivative calculus and machine learning." *Remote Sens.* 12 (13): 2082. <https://doi.org/10.3390/rs12132082>.

Bon-Gang, H. 2018. *Performance and improvement of green construction projects: Management strategies and innovations*. Oxford, UK: Butterworth-Heinemann.

Brakstad, F. 1992. "A comprehensive pollution survey of polychlorinated dibenzo-p-dioxins and dibenzofurans by means of principal component analysis and partial least-squares regression." *Chemosphere* 24 (12): 1885–1903. [https://doi.org/10.1016/0045-6535\(92\)90241-I](https://doi.org/10.1016/0045-6535(92)90241-I).

Breiman, L. 2001. "Random forests." *Mach. Learn.* 45 (1): 5–32. <https://doi.org/10.1023/A:1010933404324>.

Brown, C. B. 1943. "Discussion of 'Sedimentation in reservoirs.'" In Vol. 69 of *Proc., ASCE*, 1493–1500. Reston, VA: ASCE.

Cao, X. H., I. Stojkovic, and Z. Obradovic. 2016. "A robust data scaling algorithm to improve classification accuracies in biomedical data." *BMC Bioinf.* 17 (1): 1–10.

Choubin, B. H., H. Darabi, O. Rahmati, F. Sajedi-Hosseini, and B. Klove. 2018. "River suspended sediment modelling using the CART model: A comparative study of machine learning techniques." *Sci. Total Environ.* 615 (Feb): 272–281.

Daly, C., J. I. Smith, and K. V. Olson. 2015. "Mapping atmospheric moisture climatologies across the conterminous United States." *PLoS One* 10 (10): e0141140. <https://doi.org/10.1371/journal.pone.0141140>.

Demiray, B. Z., M. Sit, and I. Demir. 2021. "D-SRGAN: DEM super-resolution with generative adversarial networks." *SN Comput. Sci.* 2 (1): 1–11. <https://doi.org/10.1007/s42979-020-00442-2>.

Emmerson, R. H., S. O'Reilly-Wiese, C. Macleod, and J. Lester. 1997. "A multivariate assessment of metal distribution in inter-tidal sediments of the Blackwater Estuary, UK." *Mar. Pollut. Bull.* 34 (11): 960–968. [https://doi.org/10.1016/S0025-326X\(97\)00067-2](https://doi.org/10.1016/S0025-326X(97)00067-2).

- Gautam, A., M. Sit, and I. Demir. 2022. "Realistic river image synthesis using deep generative adversarial networks." *Front. Water* 4 (Feb): 784441. <https://doi.org/10.3389/frwa.2022.784441>.
- Géron, A. 2022. *Hands-on machine learning with Scikit-Learn, Keras, and TensorFlow*. Sebastopol, CA: O'Reilly Media.
- Goodfellow, I., Y. Bengio, and A. Courville. 2016. *Deep learning*. Cambridge, MA: MIT Press.
- Gorelick, N., M. Hancher, M. Dixon, S. Ilyushchenko, D. Thau, and R. Moore. 2017. "Google Earth engine: Planetary-scale geospatial analysis for everyone." *Remote Sens. Environ.* 202 (Mar): 18–27. <https://doi.org/10.1016/j.rse.2017.06.031>.
- Goyal, H., D. Sandeep, R. Vanu, R. Pokuri, S. Kathula, and N. Battula. 2014. "Normalization of data in data mining." *Int. J. Software Web Sci.* 10 (Jun): 32–33.
- Gu, X., and P. Angelov. 2017. "Autonomous anomaly detection." In *Proc., 2017 Evolving and Adaptive Intelligent Systems (EAIS)*, 1–8. New York: IEEE.
- Harrell, F. E. 2001. *Regression modeling strategies: With application to linear models, logistic regression, and survival analysis*. New York: Springer.
- Hassan, S., N. Shaukat, A. Ahmad, M. Abid, A. Hashmi, and M. Shahid. 2022. "Prediction of the amount of sediment deposition in Tarbela Reservoir using machine learning approaches." *Water* 14 (19): 3098. <https://doi.org/10.3390/w14193098>.
- Hazarika, B. B., D. Gupta, and M. Berlin. 2020. "Modeling suspended sediment load in a river using extreme learning machine and twin support vector regression with wavelet conjunction." *Environ. Earth Sci.* 79 (10): 234.
- Hecht, R. M., and N. Tishby. 2005. "Extraction of relevant speech features using the information bottleneck method." In *Proc., Interspeech*, 353–356. Wadern, Germany: Schloss Dagstuhl.
- Idrees, M., M. Jehanzaib, D. Kim, and T.-W. Kim. 2021. "Comprehensive evaluation of machine learning models for suspended sediment load inflow prediction in a reservoir." *Stochastic Environ. Res. Risk Assess.* 35 (9): 1805–1823. <https://doi.org/10.1007/s00477-021-01982-6>.
- Jirachan, T., and K. Piromsopa. 2015. "Applying KSE-test and K-means clustering towards scalable unsupervised intrusion detection." In *Proc., 2015 12th Int. Joint Conf. on Computer Science and Software Engineering (JCSSE)*, 82–87. New York: IEEE.
- Jothiprakash, V., and V. Garg. 2009. "Reservoir sedimentation estimation using artificial neural network." *J. Hydrol. Eng.* 14 (9): 1035–1040. [https://doi.org/10.1061/\(ASCE\)HE.1943-5584.0000075](https://doi.org/10.1061/(ASCE)HE.1943-5584.0000075).
- Lajczak, A. 1996. "Modelling the long-term course of non-flushed reservoir sedimentation and estimating the life of dams." *Earth Surf. Processes Landforms* 21 (12): 1091–1107. [https://doi.org/10.1002/\(SICI\)1096-9837\(199612\)21:12<1091::AID-ESP653>3.0.CO;2-2](https://doi.org/10.1002/(SICI)1096-9837(199612)21:12<1091::AID-ESP653>3.0.CO;2-2).
- Legates, D. R., and G. J. McCabe Jr. 1999. "Evaluating the use of 'goodness-of-fit' measures in hydrologic and hydroclimatic model validation." *Water Resour. Res.* 35 (1): 233–241. <https://doi.org/10.1029/1998WR900018>.
- Lehner, B., et al. 2011. "High-resolution mapping of the world's reservoirs and dams for sustainable river-flow management." *Front. Ecol. Environ.* 9 (9): 494–502. <https://doi.org/10.1890/100125>.
- Maimaitijiang, M., V. Sagan, P. Sidike, S. Hartling, F. Esposito, and F. Fritsch. 2020. "Soybean yield prediction from UAV using multimodal data fusion and deep learning." *Remote Sens. Environ.* 237 (Feb): 111599. <https://doi.org/10.1016/j.rse.2019.111599>.
- Malik, A., A. Kumar, O. Kisi, and J. Shiri. 2019. "Evaluating the performance of four different heuristic approaches with gamma test for daily suspended sediment concentration modeling." *Environ. Sci. Pollut. Res.* 26 (22): 22670–22687. <https://doi.org/10.1007/s11356-019-05553-9>.
- Manikanta, C., and V. Mamatha Jadav. 2015. "Evaluation of modified PLS regression method to fill the missing values in training dataset." In *Proc., Int. Conf. on Smart Sensors and Systems (IC-SSS)*, 1–5. New York: IEEE.
- Montavon, G., S. Wojciech, and M. Klaus-Robert. 2018. "Methods for interpreting and understanding deep neural networks." *Digital Signal Process.* 73 (Aug): 1–15. <https://doi.org/10.1016/j.dsp.2017.10.011>.
- Moriasi, D. N., J. Arnold, M. Van Liew, R. Bingner, R. Harmel, and T. Veith. 2007. "Model evaluation guidelines for systematic quantification of accuracy in watershed simulations." *Trans. ASABE* 50 (3): 885–900. <https://doi.org/10.13031/2013.23153>.
- Morris, G., and J. Fan. 1998. *Reservoir sedimentation handbook: Design and management of dams, reservoirs, and watershed for sustainable use*. New York: McGraw-Hill.
- Noble, W. S. 2006. "What is a support vector machine?" *Nat. Biotechnol.* 24 (12): 1565–1567.
- Patro, S., and K. Sahu. 2015. "Normalization: A preprocessing stage." Preprint, submitted March 19, 2015. <https://arxiv.org/abs/1503.06462>.
- Peterson, K. T., V. Sagan, P. Sidike, E. A. Hasenmueller, J. J. Sloan, and J. H. Knauft. 2019. "Machine learning-based ensemble prediction of water-quality variables using feature-level and decision-level fusion with proximal remote sensing." *Photogramm. Eng. Remote Sens.* 85 (4): 269–280. <https://doi.org/10.14358/PERS.85.4.269>.
- Peterson, K. T., V. Sagan, and J. J. Sloan. 2020. "Deep learning-based water quality estimation and anomaly detection using Landsat-8/Sentinel-2 virtual constellation and cloud computing." *GISci. Remote Sens.* 57 (4): 510–525. <https://doi.org/10.1080/15481603.2020.1738061>.
- Pinson, A., B. Baker, P. Boyd, R. Grandpre, K. White, and M. Jonas. 2016. *US Army Corps of Engineers reservoir sedimentation in the context of climate change*. Civil Works Technical Rep. (CWTS) 2016-05. Washington, DC: USACE.
- Randle, T., G. Morris, D. Tullis, F. Weirich, G. Kondolf, D. Moriasi, and D. Wegner. 2021. "Sustaining United States reservoir storage capacity: Need for a new paradigm." *J. Hydrol.* 602 (Dec): 126686. <https://doi.org/10.1016/j.jhydrol.2021.126686>.
- Renard, K. G. 1997. *Predicting soil erosion by water: A guide to conservation planning with the Revised Universal Soil Loss Equation (RUSLE)*. Washington, DC: USDA.
- Rowan, J. S., L. E. Price, C. P. Fawcett, and P. C. Young. 2001. "Reconstructing historic reservoir sedimentation rates using data-based mechanistic modelling." *Phys. Chem. Earth Part B* 26 (1): 77–82. [https://doi.org/10.1016/S1464-1909\(01\)85018-8](https://doi.org/10.1016/S1464-1909(01)85018-8).
- Sholtes, J., C. Ubing, T. Randle, J. Fripp, D. Cenderelli, and D. Baird. 2018. "Managing infrastructure in the stream environment." *J. Am. Water Resour. Assoc.* 54 (6): 1172–1184. <https://doi.org/10.1111/1752-1688.12692>.
- Singh, A., N. Thakur, and A. Sharma. 2016. "A review of supervised machine learning algorithms." In *Proc., 3rd Int. Conf. on Computing for Sustainable Global Development (INDIACom)*, 1310–1315. New York: IEEE.
- Sit, M., B. Demiray, and I. Demir. 2021a. "Short-term hourly streamflow prediction with graph convolutional GRU networks." Preprint, submitted July 7, 2021. <https://arxiv.org/abs/2107.07039>.
- Sit, M., B. C. Seo, and I. Demir. 2021b. "Towarain: A statewide rain event dataset based on weather radars and quantitative precipitation estimation." Preprint, submitted July 7, 2021. <https://arxiv.org/abs/2107.03432>.
- Sundborg, Å. 1992. "Lake and reservoir sedimentation prediction and interpretation." *Geogr. Ann.: Ser. A Phys. Geogr.* 74 (2–3): 93–100. <https://doi.org/10.1080/04353676.1992.11880353>.
- Tarela, P. A., and A. N. Menéndez. 1999. "A model to predict reservoir sedimentation." *Lakes Reservoirs: Res. Manage.* 4 (3–4): 121–133. <https://doi.org/10.1046/j.1440-1770.1999.00087.x>.
- Taylor, R. 1990. "Interpretation of the correlation-coefficient—A basic review." *J. Diagn. Med. Sonogr.* 6 (1): 35–39. <https://doi.org/10.1177/875647939000600106>.
- Tishby, N., F. Pereira, and W. Bialek. 2000. "The information bottleneck method." Preprint, submitted April 24, 2000. <https://arxiv.org/abs/0004057>.
- Trimble, S. W. 1999. "Decreased rates of alluvial sediment storage in the Coon Creek Basin, Wisconsin, 1975–93." *Science* 285 (5431): 1244–1246. <https://doi.org/10.1126/science.285.5431.1244>.
- USBR (US Bureau of Reclamation). 2006. *Hydrology, hydraulics, and sediment studies for the Matilija Dam Ecosystem Restoration Project, Venture, CA*. Draft Rep. Denver: Sedimentation and River Hydraulics Group.
- USDA SCS (USDA Soil Conservation Service). 1986. *Technical release 55: Urban hydrology for small watersheds*. 2nd ed. Washington, DC: USDA SCS.

- USGS. 2017. *1/3rd arc-second digital elevation models (DEMs)—USGS National Map 3DEP downloadable data collection*. Washington, DC: USGS.
- Verstraeten, G., J. Poesen, J. de Vente, and X. Koninckx. 2003. "Sediment yield variability in Spain: A quantitative and semiquantitative analysis using reservoir sedimentation rates." *Geomorphology* 50 (4): 327–348. [https://doi.org/10.1016/S0169-555X\(02\)00220-9](https://doi.org/10.1016/S0169-555X(02)00220-9).
- Viger, R. J., and A. Bock. 2014. *GIS features of the geospatial fabric for nation hydrologic modeling*. Washington, DC: USGS.
- Vörösmarty, C., M. Meybeck, B. Fekete, K. Sharma, P. Green, and J. Syvitski. 2003. "Anthropogenic sediment retention: Major global impact from registered river impoundments." *Glob. Planet. Change* 39 (1–2): 169–190.
- Wang, Z. Y., and C. Hu. 2009. "Strategies for managing reservoir sedimentation." *Int. J. Sediment Res.* 24 (4): 369–384. [https://doi.org/10.1016/S1001-6279\(10\)60011-X](https://doi.org/10.1016/S1001-6279(10)60011-X).
- Xiang, Z., and I. Demir. 2020. "Distributed long-term hourly streamflow prediction using deep learning—A case study for State of Iowa." *Environ. Modell. Software* 131 (Apr): 104761. <https://doi.org/10.1016/j.envsoft.2020.104761>.
- Xu, H., I. Demir, C. Koylu, and M. Muste. 2019. "A web-based geovisual analytics platform for identifying potential contributors to culvert sedimentation." *Sci. Total Environ.* 692 (Mar): 806–817. <https://doi.org/10.1016/j.scitotenv.2019.07.157>.
- Zdaniuk, B. 2014. "Ordinary least squares (OLS) model." In *Encyclopedia of quality of life and well-being research*. New York: Springer. https://doi.org/10.1007/978-94-007-0753-5_2008.
- Zounemat-Kermani, M., O. Kisi, J. Piri, and A. Mahdavi-Meymand. 2019. "Assessment of artificial intelligence-based models and metaheuristic algorithms in modeling evaporation." *J. Hydrol. Eng.* 24 (10): 04019033. [https://doi.org/10.1061/\(ASCE\)HE.1943-5584.0001835](https://doi.org/10.1061/(ASCE)HE.1943-5584.0001835).
- Zounemat-Kermani, M., A. Mahdavi-Meymand, M. Alizamir, S. Adarsh, and Z. Yaseen. 2020. "On the complexities of sediment load modeling using integrative machine learning: Application of the great river of Loiza in Puerto Rico." *J. Hydrol.* 585 (Aug): 124759. <https://doi.org/10.1016/j.jhydrol.2020.124759>.