

COMPARATIVE ANALYSIS ON SNOWMELT-DRIVEN STREAMFLOW FORECASTING USING MACHINE LEARNING TECHNIQUES

A PREPRINT

UKESH THAPA

AI Research Center
Advanced College of Engineering and Management,
Tribhuvan University
Kathmandu, Nepal
ukesh.thapa@acem.edu.np

Bipun Man Pati

AI Research Center
Advanced College of Engineering and Management,
Tribhuvan University
Kathmandu, Nepal
bipunmanpati@acem.edu.np

Samit Thapa

Department of Civil Engineering
Advanced College of Engineering and Management,
Tribhuvan University
Kathmandu, Nepal
samit.thapa@acem.edu.np

Dhiraj Pyakurel

Department of Electronics and Computer Engineering
Advanced College of Engineering and Management,
Tribhuvan University
Kathmandu, Nepal
dhiraj@acem.edu.np

Anup Shrestha

Department of Electronics and Computer Engineering
National College of Engineering,
Tribhuvan University
Lalitpur, Nepal
anup@nce.edu.np

April 24, 2024

ABSTRACT

The rapid advancement of machine learning techniques has led to their widespread application in various domains including water resources. However, snowmelt modeling remains an area that has not been extensively explored. In this study, we propose a state-of-the-art (SOTA) deep learning sequential model, leveraging the Temporal Convolutional Network (TCN), for snowmelt-driven discharge modeling in the Himalayan basin of the Hindu Kush Himalayan Region. To evaluate the performance of our proposed model, we conducted a comparative analysis with other popular models including Support Vector Regression (SVR), Long Short Term Memory (LSTM), and Transformer. Furthermore, Nested cross-validation (CV) is used with five outer folds and three inner folds, and hyper-parameter tuning is performed on the inner folds. To evaluate the performance of the model mean absolute error (MAE), root mean square error (RMSE), R square (R^2), Kling-Gupta Efficiency (KGE), and Nash-Sutcliffe Efficiency (NSE) are computed for each outer fold. The average metrics revealed that TCN outperformed the other models, with an average MAE of 0.011, RMSE of 0.023, R^2 of 0.991, KGE of 0.992, and NSE of 0.991. The findings of this study demonstrate the effectiveness of the deep learning model as compared to traditional machine learning approaches for snowmelt-driven streamflow forecasting. Moreover, the superior performance of TCN highlights its potential as a promising deep learning model for similar hydrological applications.

Keywords Support Vector Regression (SVR) · Long Short Term Memory (LSTM) · Transformer · Temporal Convolutional Network (TCN) · Nested Cross-Validation · Snowmelt

1 Introduction

1.1 Hindu Kush Himalayan Region

Snowmelt is the process by which the snow or ice on the ground or various surfaces changes into water due to the temperature rise. Snowmelt plays a pivotal role in the environment and human society development. Snowmelt is also a major source of fresh water in the world. Central Asia contains a large amount of ice in the Hindu Kush-Himalayan (HKH) region and is highly sensitive to global climate change, facing significant warming ($0.21 \pm 0.08^\circ/\text{decade}$) over the past few decades. In Panday et al. [2011] a dynamic threshold-based method was applied and observation was done from (2000-2008), where average melt duration was calculated and longer melt season (~ 5 weeks) occurred in the eastern Himalayan region relative to central, western Himalayan, and the Karakoram region. People residing in the river basins originating from the snow-dominated HKH region rely on rivers for food, water supply, and electricity (Wester et al. [2019]). The seasonal snowmelt process significantly affects the ecosystem, water availability, agriculture, and water resources. Additionally, the snowmelt directly influences the river flow, aquatic habitats, and the functioning of the hydroelectric power system. Therefore, precise forecasting of the snowmelt runoff is essential in this area for efficient planning and management. The geographical conditions and extreme weather in this region lead to inadequate ground-truth information for developing the optimum water resource utilization strategy. Therefore, remotely sensed snow and meteorological data are indispensable tools for traditional ground-based monitoring.

1.2 Snowmelt Forecasting Using Machine Learning

Machine learning (ML) has diverse applications across interdisciplinary fields, showcasing its versatility and impact. The application of ML in the domain of water resources has been well investigated in previous studies (on Application of Artificial Neural Networks in Hydrology [2000a,b]). Callegari et al. [2015] proposes Support Vector Regression (SVR) for snowmelt-driven discharge forecasting in the Italian Alps using Snow Cover Area (SCA), antecedent discharge (Q), and metrological data, which has outperformed a simple linear auto-regressive model with an average of 33% relative root mean square error (RMSE) on test samples. De Gregorio et al. [2018] demonstrates the applicability of the SVR model in operational river discharge forecasting, where the SVR model performed better on a single gauging station with a mean improvement of about 48% in the RMSE. In a study, Uysal et al. [2016] used a simple Artificial Neural Network (ANN) for snowmelt runoff prediction, where SCA, Q, temperature (T), and precipitation (P) were used as inputs to the model, and research verified that the ANN model outperformed any temperature index (TI) model as its Nash-Sutcliffe model efficiency (NSE) increased from 0.51 to 0.71 in forecasting. In a study by Najafzadeh et al. [2023], an improved version of the robust ML models is proposed to compare the performance of the gene-expression programming (GEP) ($R = 0.6089$ and $\text{RMSE} = 7.7229$), evolutionary polynomial regression (EPR) ($R = 0.6020$ and $\text{RMSE} = 1.6923$), multivariate adaptive regression spline (MARS) ($R = 0.4333$ and $\text{RMSE} = 11.5046$), and model tree (MT) to find the ecological status of the river. In Najafzadeh and Anvari [2023] has used three different ML techniques such as GEP, MT, and MARS for streamflow forecasting; they claimed that MT had lower uncertainty ($95\% \text{PPU} = 0.59$ and $R\text{-factor} = 1.67$) than other models.

In past research studies, traditional ANNs were used for the time series problem, but it was not enough to capture all the necessary or complex information from the data. The advancement of technology leads to the development of more complicated architectures. At first, a study proposed by Nagesh Kumar et al. [2004] used a Recurrent Neural Network (RNN) for the time series problem, but due to its drawbacks, like the exploding and vanishing gradient problems, Long Short-Term Memory (LSTM) overcame all of those problems. In their research, Kratzert et al. [2018] used two-layered LSTM for rainfall-runoff modeling with NSE 0.63 and argued that it can also mimic the snowmelt process by learning the relationship between precipitation during winter and runoff in spring, but they did not investigate the influence of the hyper-parameter tuning on the model performance. Le et al. [2019] proposes an LSTM model that has an average error of $150 \text{ m}^3/\text{s}$ RMSE with NSE above 99% while forecasting discharge for one day. Le et al. [2019] and Fan et al. [2020] have claimed that window size in forecasting problems is an important hyper-parameter to be tuned for the best model performance, but in both studies, other hyper-parameters, such as the number of LSTM layers and optimizers, were not observed. In Granata et al. [2022], have done a comparative analysis of random forest (RF), SVR, and Multi-layered perceptron (MLP) for the estimation of daily volumetric soil water content, where MLP has the best performance with Coefficient of determination (R^2) value 0.957. Similarly, In Najafzadeh et al. [2024] a comparative analysis of nine different ML models is done to predict the vulnerability of the rip current event using two parameters such as dimensionless fall velocity parameter (Ω) and tide range (TR). In a recent study, Thapa et al. [2020a] has compared snowmelt prediction on different models by hyper-tuning parameters with the hit and trail method and setting window size for prediction. Also, gamma testing was done where SCA, T, and Q were chosen as inputs for different models, out of which the LSTM model has better performance with 0.997, 0.112, 0.173, 0.99, and 0.995 as R^2 , mean absolute error (MAE), RMSE, Kling-Gupta efficiency (KGE), and NSE respectively. This study did not evaluate the performance of

the model using nested cross-validation (CV), the state-of-the-art (SOTA) model created only performed better for the specified dataset but not for any unseen dataset.

The main objective of this study is to investigate the SOTA deep learning (DL) architectures such as the Transformer and the Temporal Convolutional Networks (TCN) to perform snowmelt prediction over traditional ML i.e. SVR and old DL techniques i.e. LSTM. In the snowmelt prediction, the relationship between the input and the target variable can be complex and dynamic, so we compare the performance of the traditional ML and DL techniques to handle the complex pattern of the data for forecasting. In the current literature reviews, we found no consistent metrics are being used to compare model improvement. We have compared the four different models and have used consistent metrics (such as KGE, NSE, RMSE, and MAE) to create a benchmark for the model to compare its performance. Furthermore, we evaluate the performance of each model using nested CV for the model generability test. In this study, SCA, T, Q, and remotely sensed daily P are used as inputs. The remainder of the paper is organized as follows: Section 2 provides the materials and methods of the study. Section 3 then describes the methodology. The results of the proposed study are then discussed in Section 4. The discussion part is discussed in Section 5. Finally, Section 6 provides the conclusion and directions for future research.

2 Materials and Methods

2.1 Study Area

In the Langtang basin, which is situated in Nepal's Central Himalayas, the study's main focus is on predicting the runoff from snowfall. Glaciers occupy 110 km² of the 354 km² study area, which spans from 3647 to 7213 meters above sea level. RGI-GLIMS version 6 calculates the extent of glaciers (Consortium et al. [2017]). The entire amount of water discharge in this watershed is largely determined by the melting of glaciers and snow (Ragettli et al. [2015]). Figure 1 shows the HKH area which is used for the dataset collection.

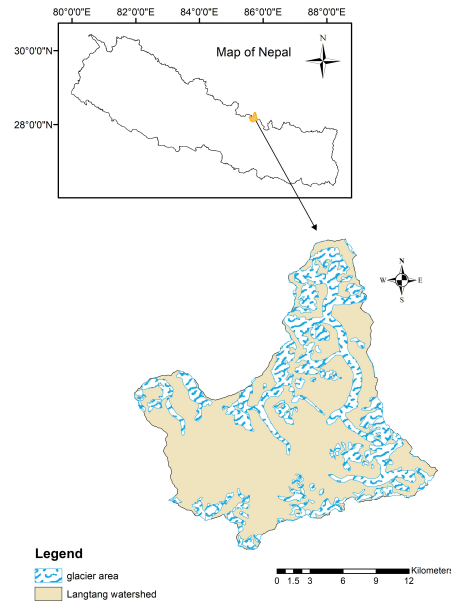


Figure 1: Map of Langtang basin area

2.2 Hydrometeorological Data

In Nepal, the Department of Hydrology and Meteorology handed the hydrological data from the Kyangjing station (28.216° latitude, 85.55° longitude). The APHRODITE product (APHRO_TAVE_MA_V1808) handed diurnal temperature data with a $0.25^\circ \times 0.25^\circ$ grid for the times 2002–2012 [Yasutomi et al. [2011]]. Previous exploration in the Langtang [Thapa et al. [2020b]] receptacle revealed a high correlation (Spearman's $\rho > 0.9$) between APHRODITE products and ground compliances. The rush data were attained using 3B42RT TRMM lines, which were sourced from

the Tropical Rainfall Measuring Mission (TRMM) at a resolution of 0.25 degrees. These datasets combine information from multiple detectors including satellite infrared data to estimate rain. The 3B42RT TRMM dataset is created by combining data from rain gauges. TRMM products are available at <https://pmm.nasa.gov> and are constantly used in the Himalayan region (Immerzeel et al. [2009]).

2.3 Snow Cover

The MODIS MOD10A2 interpretation six products are employed to collude snow cover. Stigter et al. [2017], has observed an accuracy of 83.1 % with on-point snow compliances in the Langtang receptacle, the snow mapping algorithm utilizes MODIS bands 4 and 6 to calculate the regularized-difference snow indicator (Hall et al. [2002]). In MOD10A2, a pixel is marked as snow if observed at least once in eight days, as no snow if absent in all eight days, and as pall if pall cover is observed every day. For this study, 496 images from 2002 to 2012 were obtained from the National Snow and Ice Data Center’s (NSIDC) website. The 30 m-resolution ASTER Global Digital Elevation Model (GDEM) was sourced from <https://lpdaac.usgs.gov>. The study area boundary was outlined in the DEM. MOD10A2 images are projected to the World Geodetic System 1984 (WGS84), UTM Zone 45. The 8-day maximum SCA was extracted from projected snow images using the delineated boundary. Images with more than 10% cloud cover were discarded. Finally, the daily SCA for 365 days in a year was interpolated or extrapolated from the 8-day maximum SCA using the cubical spline method.

2.4 Dataset Preparation

SCA, Q, T, and P are used as inputs for forecasting the snowmelt runoff. MinMaxScaler scales the data without losing its distribution for faster convergence during the training process. MinMaxScaler uses data value and subtracts it from the minimum value of the overall features which is then divided by the difference of the maximum value and the minimum value of features. We have used two days’ data to predict the third day’s SCA i.e. window size of 2. Furthermore, the nested CV of the model is carried out to check the generability of any unseen data with five folds for the outer loop and three folds for the inner loop. The tensor shape for our models is (window size, number of features) i.e. (2,3) for three inputs and (2,4) for four inputs. Keras’ K-Fold is used for splitting the dataset with a random state of 42 with a splitting ratio of 80% i.e. 3212 samples and 20% i.e. 803 samples for training and testing respectively.

$$MinMaxScaler(m) = \frac{x - x_{min}}{x_{max} - x_{min}}, \quad (1)$$

where,

x is the input value,

xmin is the minimum value of the column,

xmax is the maximum value of the column

2.5 Experimental setup

The simulations are performed for four different architectures:

1. SVR
2. LSTM
3. Transformer
4. TCN

Each of the four different architectures was evaluated for two different inputs, M1 with inputs (T, Q, SCA, and P) and M2 with (T, Q, and SCA), respectively. The list of hyper-parameters used for the experiments are given in the table 1.

Performance of all the four architectures are evaluated using the following metrics:

1. KGE: KGE provides an overview of bias ratio, correlation, and variability. The value of KGE equal to 1 indicates a perfect agreement between simulations and observations.

$$KGE = 1 - \sqrt{(r - 1)^2 + (\alpha - 1)^2 + (\beta - 1)^2}, \quad (2)$$

where,

$\beta = \frac{\bar{Q}'}{\bar{Q}}$ is the bias ratio,

$\bar{Q}$ is the average observed discharge,

$\bar{Q}'$ is the average simulated discharge,

$\gamma = \frac{CV_s}{CV_o}$ is the variability,

CV_o is the observed coefficient of variation,

CV_s is the simulated coefficient of variation, and

r is Pearson's correlation coefficient.

The optimum value of KGE r , β , and γ is 1. No cross-correlation between the bias ratio and variability is ensured by using CV in calculating γ (Kling et al. [2012]).

2. NSE: The magnitude of residual variance concerning the observed data variance is provided by the NSE.

$$NSE = 1 - \frac{\sum_{t=1}^N (Q_t - Q'_t)^2}{\sum_{t=1}^N (Q_t - \bar{Q})^2}, \quad (3)$$

where,

Q_t is the observed discharge at time t ,

Q'_t is the simulated discharge at time t , and

$\bar{Q}$ denotes average observed discharge.

3. R^2 : The coefficient of determination (R^2) provides the intensity of the relation between the measured and the predicted values. Its value ranges from 0 to 1, closer to 0 is a low correlation whereas, 1 represents a high correlation.

$$R^2 = \left[\frac{\sum_{t=1}^n (Q'_t - \bar{Q}')(Q_t - \bar{Q})}{\sqrt{\sum_{t=1}^n (Q'_t - \bar{Q}')^2} \sqrt{\sum_{t=1}^n (Q_t - \bar{Q})^2}} \right]^2, \quad (4)$$

where

Q_t is the observed discharge at time t ,

Q'_t is the simulated discharge at time t ,

$\bar{Q}$ is the average observed discharge, and

$\bar{Q}'$ is the average simulated discharge.

4. RMSE: RMSE is an evaluation metric used to measure the average differences between predicted and actual values. It represents the standard deviation of the error, with a lower RMSE value indicating a better fit. RMSE is sensitive to large errors and is calculated using the following equation:

$$RMSE = \sqrt{\frac{\sum_{t=1}^n (Q'_t - Q_t)^2}{n}}, \quad (5)$$

where

Q_t is the observed discharge at time t , and

Q'_t is the simulated discharge at time t .

5. MAE: MAE is the absolute difference between measured and predicted values. A lower MAE represents lower error. The equation of MAE is given below:

$$MAE = \frac{1}{n} \sum_{t=1}^n |Q_t - Q'_t|, \quad (6)$$

where,

Q_t is the observed discharge at time t , and

Q'_t is the simulated discharge at time t .

Table 1: Hyper-parameters to tune four different models for optimal performance

Model	Hyper-parameter	Values
SVR	C	[0.1, 1, 10]
	Epsilon	[0.01, 0.1, 0.2]
	Kernel	[Linear, RBF]
LSTM	LSTM Layers	[1, 2, 3]
	LSTM Units	[32, 64, 128]
	Dropout Rate	[0.2, 0.3, ..., 0.5]
	Optimizer	[Adam, Adamax, RMSProp, SGD]
	Learning Rate	[0.0001, 0.001, ..., 0.1]
Transformer	Transformer Blocks	[2, 4, 6, 8]
	Head Size	[8, 16, ..., 256]
	Num Heads	[2, 4, ..., 16]
	FF Dim	[4, 8, ..., 64]
	Dropout Rate	[0.1, 0.2, ..., 0.6]
	Num MLP Layers	[1, 2, 3]
	MLP Units	[32, 64, ..., 256]
	MLP Dropout	[0.1, 0.2, ..., 0.6]
	Optimizer	Adam, Adamax, RMSProp, SGD
	Learning Rate	[0.0001, 0.001, ..., 0.1]
TCN	TCN Layers	[1, 2, 3]
	Num Filters	[32, 64, ..., 128]
	Kernel Size	[2, 3, 4]
	Optimizer	[Adam, Adamax, RMSProp, SGD]
	Learning Rate	[0.0001, 0.001, ..., 0.1]

3 Methodology

In this section, the proposed methodology for snowmelt runoff prediction is shown in Figure 5. At first, we preprocess the dataset using MinMaxScaler to scale the data and set the value of Windows size 2 for the third day’s SCA prediction. The prepared data are then split into outer and inner folds in a random state of 42 with a splitting ratio of 80% i.e. 3212 samples and 20% i.e. 803 samples for training and testing respectively. The split data are used for the traditional ML and DL architectures as shown in Figure 5. The trained models are evaluated using metrics such as RSME, MAE, R^2 , KGE, and NSE which are used for the models’ comparison. The performance metrics help to create the baseline for model performance comparison and choose the best-performing traditional ML or DL model.

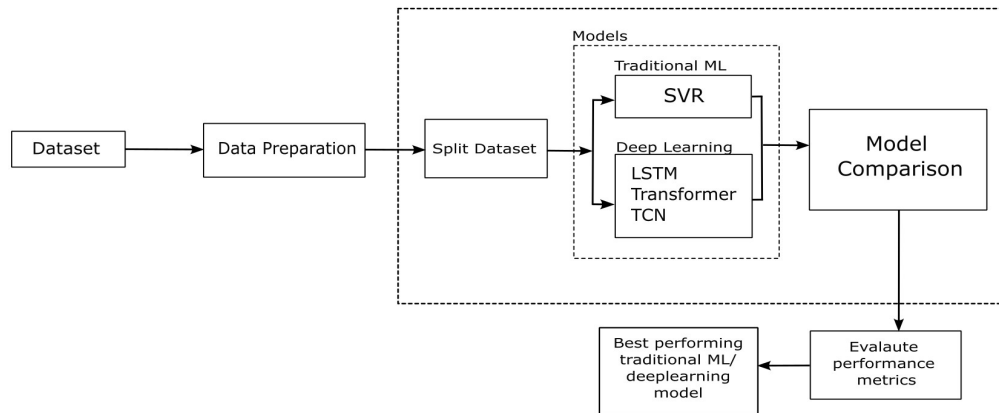


Figure 2: Methodological framework for assessment of performance of ML techniques for snowmelt forecasting

In the figure 2, the model section contains traditional ML and deep learning architecture. The detailed explanations of four different models used in the simulation are described below:

3.1 SVR

The SVR, proposed by Vapnik [1999] is a traditional ML algorithm for classification and regression. SVR can use both linear and non-linear kernels. SVR determines the hyperplane in high-dimensional space to separate the classes or output values. The main objective of SVR is to minimize the error value between the predicted and actual value for the correct prediction. In this study, we have used a Radial Basic Function (RBF) kernel with SVR, the RBF kernel is capable of capturing non-linear patterns. The rbf kernel mathematical expression is given below.

$$K(x_i, x_j) = \exp \left(-\frac{\|x_i - x_j\|^2}{2\sigma^2} \right), \quad (7)$$

where,

x_i and x_j are the input feature vectors, and

σ is a parameter controlling the width of the kernel.

In this study, we have used 0.1, 1, and 10 as regularization parameters (C), and 0.01, 0.1, and 0.2 are epsilon (ϵ) values used for the hyper-parameter tuning. We have also used both linear and RBF in the kernel, all of those hyper-parameters are determined by the grid search approach. The best hyperparameters after tuning are 0.1 for C, 0.01 for epsilon, and linear kernel.

3.2 LSTM

LSTM proposed by Hochreiter and Schmidhuber [1997] is a variation of RNN used widely in DL architecture. Unlike the vanilla RNN, the LSTM is capable of capturing the long-term dependencies and solving the problem of exploding and vanishing gradient problems. LSTM consists of feedback connections, which allow it to process sequences of data. It consists of three different gates: the input gate, the forget gate, and the output gate. The forget gate removes the information that is no longer useful in the cell state. Additional new information is added in the cell state by the input gate. The output gate extracts the useful information from the cell state to be presented as output. A classic LSTM cell is shown in figure 3 and its related equations are below.

Forget gate: $f_t = \sigma(w_f \cdot [h_{t-1}, x_t] + b_f)$

Input gate: $i_t = \sigma(w_i \cdot [h_{t-1}, x_t] + b_i)$

Update vector: $\hat{c}_t = \tanh(w_c \cdot [h_{t-1}, x_t] + b_c)$

Cell state: $f_t \odot C_{t-1} + i_t \odot \hat{c}_t$

Output gate: $o_t = \sigma(w_o \cdot [h_{t-1}, x_t] + b_o)$

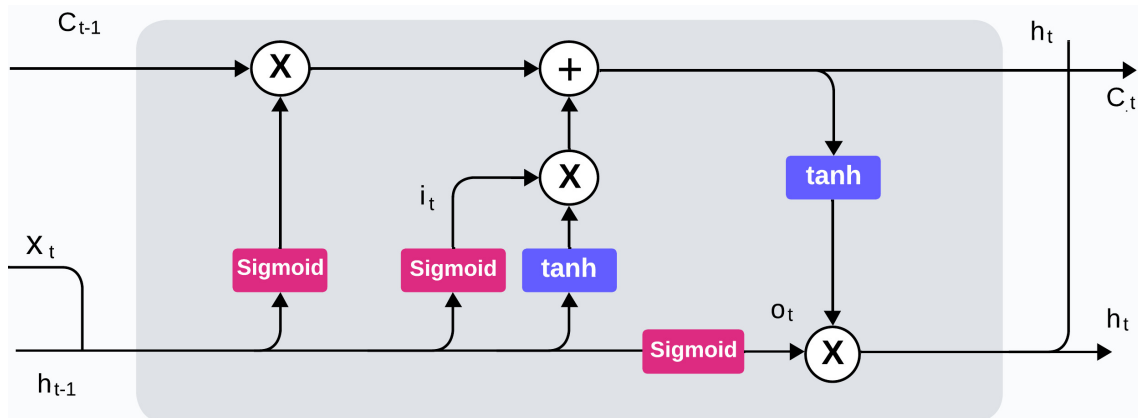


Figure 3: A classic LSTM cell, where *sigmoid* denotes sigmoidal function, *tanh* denotes a hyperbolic tangent function, *C* denotes cell state, *h* denotes hidden state, *o* denote output gate, *i* denote input vectors at time step *t*

where,

w_f , w_i , w_c , and w_o are weights for forget gate, input gate, cell state, and output gate respectively,

b_f , b_i , b_c , and b_o are biases for the forget gate, input gate, cell state, and output gate respectively,

σ represents the sigmoid activation function, and

$[h_{t-1}, x_t]$ represents the concatenation of the current input and the previous hidden state. The hyper-parameters used in the LSTM model are three different numbers of LSTM layers, three different LSTM hidden units with values 32, 64, and 128, dropout values starting from 0.2 to 0.5, learning rate values starting from 0.0001 to 0.1, and five different optimizers such as Adam, Adamax, RMSProp, and SGD are used for tuning model. The hyper-parameter tuning is done using a random search in the Keras tuner. For the optimal hyperparameters in our model, we utilized an Adam optimizer with a learning rate of 0.004, dropout rate of 0.2, and incorporated two LSTM layers each consisting of 96 hidden units.

3.3 Transformer

A transformer proposed by Vaswani et al. [2017] is a DL architecture, that includes an encoder and decoder. These encoder and decoder are connected through the attention mechanism. This architecture requires less training time than LSTM, as its latest version is highly used for training large language models (LLM). This architecture is applicable for natural language processing (NLP) and computer vision, but also in audio and multi-modal processing. Figure 4 shows transformer model architecture and its components are described below.

1. **Self-Attention Mechanism:** The self-attention mechanism is responsible for capturing any crucial information from the long sequence of data. The calculation for attention can be expressed as a large matrix calculation using the softmax function.

$$\text{Attention}(\mathbf{Q}, \mathbf{K}, \mathbf{V}) = \text{softmax} \left(\frac{\mathbf{Q}\mathbf{K}^T}{\sqrt{d_k}} \right) \mathbf{V}, \quad (8)$$

where,

$\mathbf{Q}$ represents Query,

$\mathbf{K}$ represents key,

$\mathbf{V}$ represents value, and

d_k represents the dimensionality of the key in items

2. **Multi-Head Attention:** The multi-head attention is a feature of the transformer that helps the model process different aspects of the input sequence simultaneously. In this process, multiple self-attention heads work in parallel, and outputs are concatenated and linearly transformed.

$$\text{MultiHead}(\mathbf{Q}, \mathbf{K}, \mathbf{V}) = \text{Concat}(\mathbf{h}_1, \dots, \mathbf{h}_n) \mathbf{W}_o, \quad (9)$$

where,

$\mathbf{h}$ represents the heads, and

$\mathbf{W}_o$ represents the matrix of the entire multi-head attention.

3. **Position Encoding:** Position encoding in the architecture is responsible for the order of the words in a sequence, and this is added to the input embedding.

$$\text{PE}(\text{pos}, 2i) = \sin \left(\frac{\text{pos}}{1000^{(2i/d_{\text{model}})}} \right), \quad (10)$$

$$\text{PE}(\text{pos}, 2i + 1) = \cos \left(\frac{\text{pos}}{1000^{(2i/d_{\text{model}})}} \right), \quad (11)$$

where,

pos is the position of the sequence, and

d_{model} is the dimensionality of the model.

In the transformer, we use four different transformer blocks with values 2,4,6, and 8, head size starting from 8 to 256, number of heads from 2 to 16, dropout rate with values starting from 0.1 to 0.6, three different numbers of multi-layer perceptron (MLP) layers, MLP units values from 32 to 256, MLP Dropout values starting from 0.1 to 0.6, four different optimizers such as Adam, SGD, RMSProp, and Adamax, and learning rate with values from 0.0001 to 0.1 are used as hyper-parameters to tune transformer model. We have used a Keras tuner for the hyper-parameter tuning. The optimal

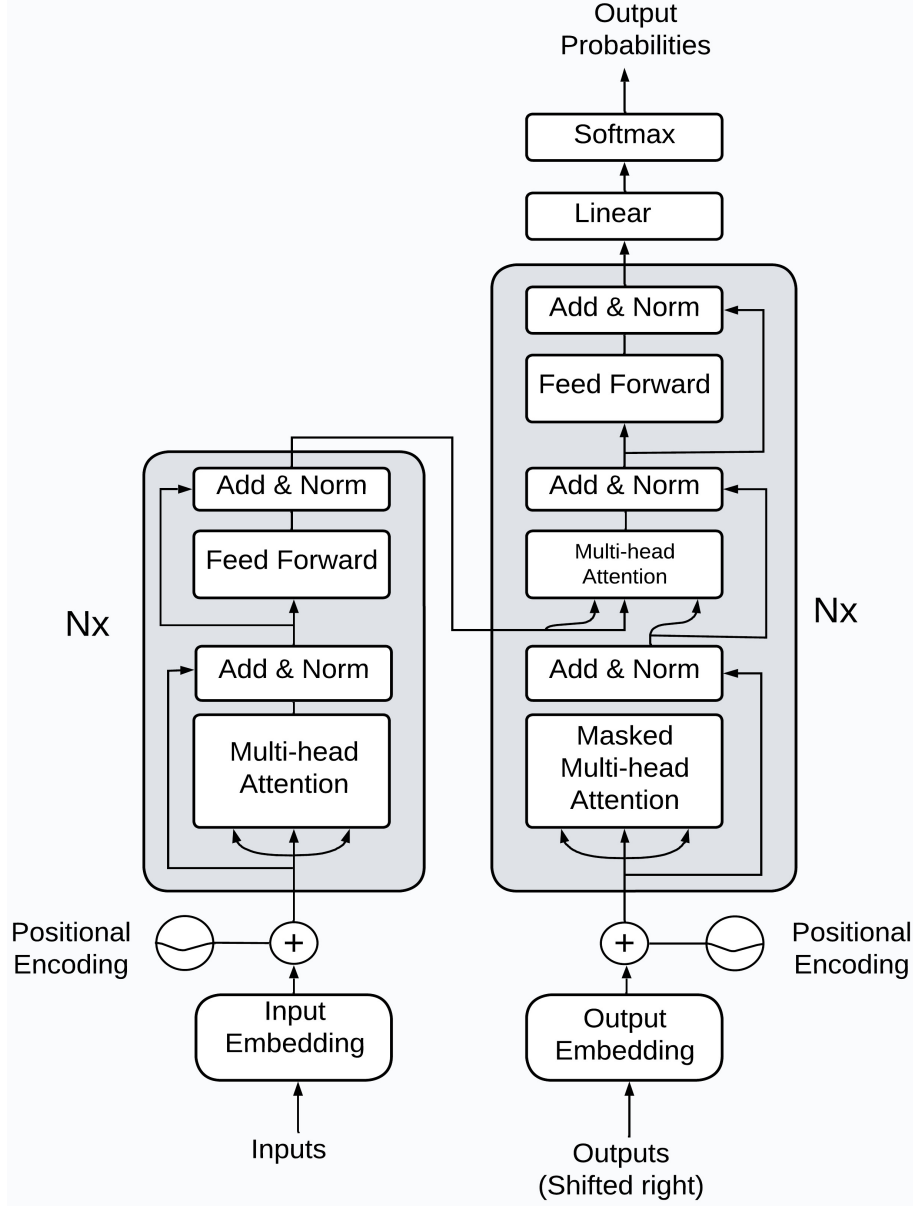


Figure 4: The detail model architecture of the Transformer. [Vaswani et al. [2017]]

hyperparameters are two transformer blocks, each with a head size of 136, and two attention heads. The two MLP layers were employed with units of 192 and 160, respectively. Dropout regularization was applied across the model, with rates of 0.41 for the first layer, 0.10 for the second layer, and 0.11 for the subsequent layers. The Adam optimizer with a learning rate of 0.0044 was used for training.

3.4 TCN

The TCN proposed by Bai et al. [2018] is a DL architecture that is capable of capturing long-term dependencies and temporal patterns. TCN architecture is capable of preventing data leakage using causal convolution by ensuring that the model's prediction is solely on past and present information. The TCN architecture can handle the sequence of varying lengths, allowing for efficient capture of long-range dependencies. Figure 5 shows the TCN model architecture and its key components are described below.

1. Dilated Causal Convolutions: In TCN, the dilated causal convolutions can process any length of sequence at varying receptive field sizes. The dilated convolution is responsible for capturing the long dependencies and is useful for memory management.

$$y_t = \sum_{i=1}^k w_i \cdot x_{t-d_i}, \quad (12)$$

where,

y_t is the output at time t ,

w_i is the weight of the convolutional filter,

x_t is the input at time t ,

d is the dilation rate, and

k is the filter size.

2. Residual Blocks: TCN uses residual connections within its blocks which help to solve the vanishing gradient problem and allow for effective optimization of deeper architecture. TCN's residual block consists of two layers of dilated causal convolution and non-linearity, for which a rectified linear unit is being used.
3. Temporal Skip connection: TCN has skip connections that connect every other layer, which makes the network capable of capturing and reusing the information. The skip connection is also useful for promoting gradient flow also can enhance the model's ability to learn hierarchical representation.

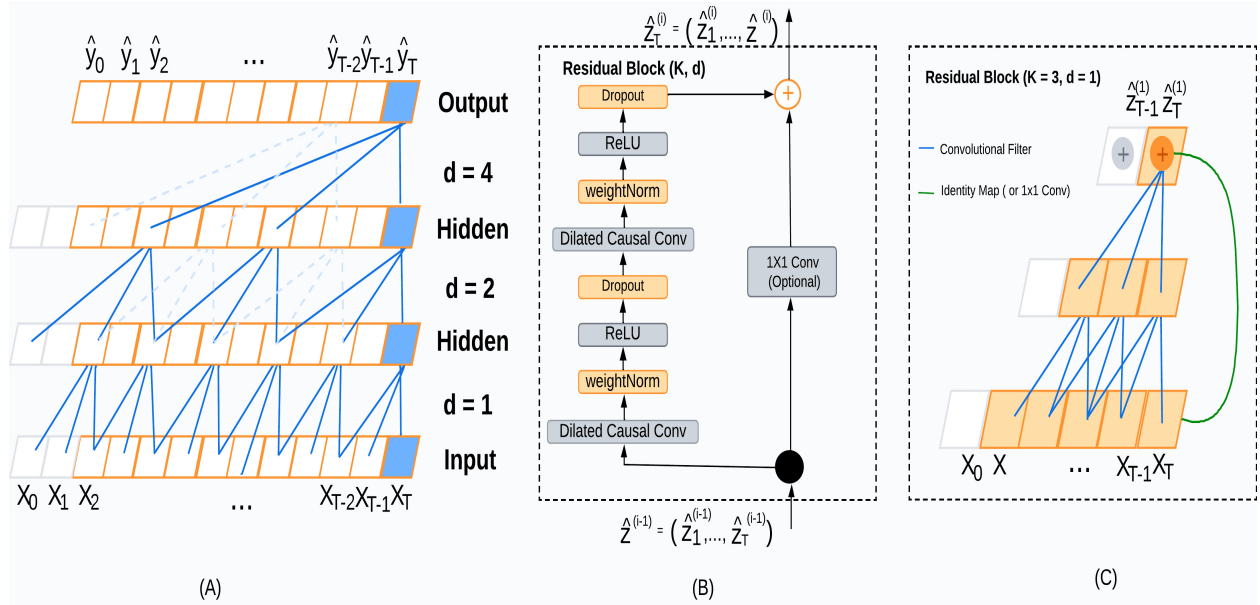


Figure 5: A TCN architectural elements. (A) A dilated causal convolution with dilated factors $d = 1, 2, 4$ and filter size $k = 3$. (B) TCN residual block. A 1×1 convolution matrix is added when residual input and output have different dimensions. (C) A residual connection example is in a TCN where the blue line represents residual function and the green lines represent identity mappings. (Bai et al. [2018])

In this study, the TCN model uses three different TCN layers, we used 32 to 256 filters, three kernel sizes such as 2, 3, and 4, five different optimizers like Adam, Adamax, RMSProp, and SGD, and learning rates values starting from 0.0001 to 0.1 use as hyper-parameters for tuning TCN model. we have used a keras tuner for hyper-parameter tuning. The optimal hyperparameters are a single layer of TCN with 96 filters, and a kernel size of 3 with Adam as optimizer along a learning rate of 0.0057.

4 Results

This section presents a comprehensive analysis of the model's performance across different scenarios. We evaluate the effectiveness of both three-input and four-input models in each fold, comparing their performance against other models. Additionally, we provide insights into the testing time required for each model.

4.1 Performance analysis of ML models with four inputs

This subsection describes the performance of ML models with four inputs in each fold of nested CV as given in table 2.

Table 2 shows the performance of the DL architectures and the traditional ML algorithms. SVR model shows the least performance for each fold of nested CV among all other models with average MAE, RMSE, R^2 , KGE, and NSE values of 0.032, 0.042, 0.97, 0.918, and 0.97, respectively. The LSTM model shows a significant performance improvement over SVR in terms of average MAE, RMSE, R^2 , KGE, and NSE of 0.013, 0.025, 0.989, 0.975, and 0.989, respectively. The transformer achieves better performance as compared to LSTM with an average KGE value of 0.986. However, TCN shows superior performance as compared to other models due to its capability of capturing long-range dependencies where average performance metrics values of MAE, RMSE, R^2 , KGE, and NSE are 0.011, 0.023, 0.991, 0.992, and 0.991, respectively. It's important to note that we used data points from the testing set of the best-predicted model fold to avoid data leakage resulting from systematic data separation using nested CV. In the scatter plot diagram, the points that lie near the diagonal line (1:1 line), estimate the river discharge accurately, whereas points above and below the diagonal overestimate and underestimate the actual flow, respectively. Figure 7 shows that the SVR and transformer model predicts good for the medium flow but the higher and lower are overestimated and underestimated, respectively. The LSTM model predicts good for low and medium flows but the high flows are underestimated by the model except the TCN model. Therefore, the TCN model predicts peak flows accurately as compared to other models.

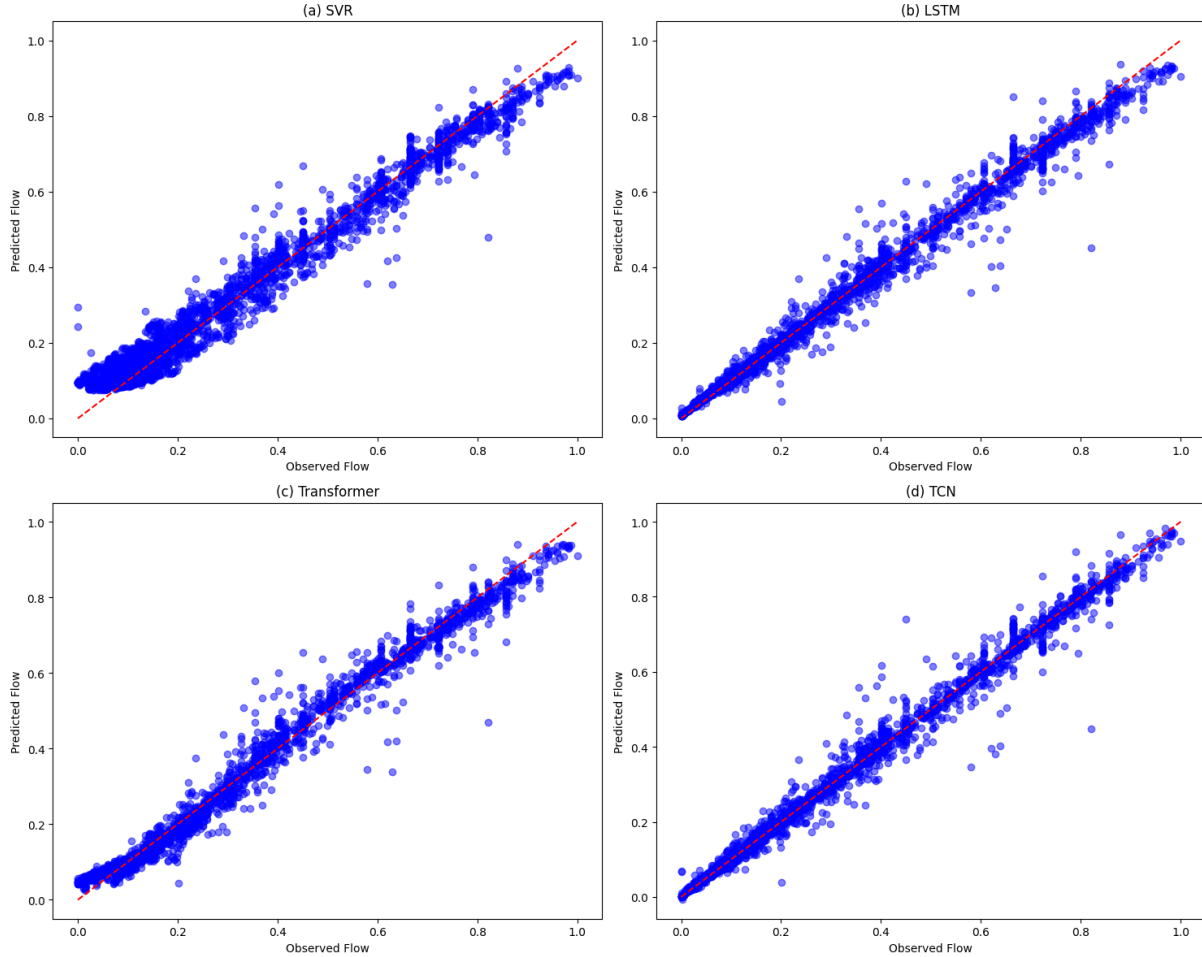


Figure 6: Scatter plot of predicted flow with observed flow (a) SVR model (b) LSTM model (c) Transformer model (d) TCN model. The solid line is the 1:1 line.

4.2 Performance analysis of ML models with three inputs

In this subsection, the performance of ML models with three inputs in each fold of nested CV is given in the table 3.

Table 2: Performance Comparison of different models with four inputs in each folds by using MAE, RMSE, R^2 , KGE, and NSE

Models	Folds	MAE	RMSE	R^2	KGE	NSE
SVR	1	0.0341	0.0451	0.9657	0.8967	0.9657
	2	0.0312	0.0401	0.9723	0.9291	0.9723
	3	0.0306	0.0406	0.9710	0.9227	0.9710
	4	0.0338	0.0450	0.9672	0.9175	0.9672
	5	0.0313	0.0395	0.9734	0.9220	0.9734
LSTM	1	0.0161	0.0301	0.9847	0.9727	0.9847
	2	0.0130	0.0235	0.9905	0.9889	0.9905
	3	0.0126	0.0239	0.9901	0.9693	0.9903
	4	0.0130	0.0242	0.9905	0.9713	0.9905
	5	0.0121	0.0220	0.9918	0.9735	0.9918
Transformer	1	0.0195	0.0308	0.9840	0.9855	0.9840
	2	0.0177	0.0274	0.9870	0.9880	0.9870
	3	0.0186	0.0283	0.9860	0.9819	0.9860
	4	0.0185	0.0281	0.9872	0.9859	0.9872
	5	0.0177	0.0258	0.9887	0.9883	0.9887
TCN	1	0.0126	0.0269	0.9878	0.9914	0.9878
	2	0.0110	0.0227	0.9911	0.9927	0.9911
	3	0.0107	0.0214	0.9919	0.9913	0.9919
	4	0.0109	0.0230	0.9914	0.9926	0.9914
	5	0.0098	0.0192	0.9937	0.9931	0.9937

SOTA DL architectures achieve higher performance as compared to the traditional ML algorithms as shown in table 3. SVR models have the minimum performance over each fold in comparison to other models with average MAE, RMSE, R^2 , KGE, and NSE values of 0.030, 0.039, 0.974, 0.937, and 0.974, respectively. The LSTM model shows a significant performance improvement over SVR with average MAE, RMSE, R^2 , KGE, and NSE of 0.012, 0.024, 0.99, 0.981, and 0.99, respectively. The LSTM shows better performance with an average KGE value of 0.981 than the transformer model with an average KGE value of 0.977. However, TCN shows superior performance also for the 3 inputs due to its capability of capturing long-range dependencies where average performance metrics values of MAE, RMSE, R^2 , KGE, and NSE are 0.011, 0.023, 0.991, 0.987, and 0.991, respectively. Figure 7 shows that the models predict good for low and medium flows but the high flows are underestimated by all models except the TCN model.

Table 3: Performance Comparison of different models with three inputs in each folds by using MAE, RMSE, R^2 , KGE, and NSE

Models	Folds	MAE	RMSE	R^2	KGE	NSE
SVR	1	0.0317	0.0426	0.9693	0.9316	0.9693
	2	0.0300	0.0386	0.9743	0.9403	0.9743
	3	0.0288	0.0374	0.9755	0.9405	0.9755
	4	0.0300	0.0395	0.9747	0.9335	0.9746
	5	0.0284	0.0368	0.9769	0.9369	0.9769
LSTM	1	0.0143	0.0272	0.9875	0.9768	0.9875
	2	0.0136	0.0236	0.9903	0.9757	0.9903
	3	0.0112	0.0231	0.9905	0.9812	0.9905
	4	0.0116	0.0232	0.9912	0.9840	0.9912
	5	0.0106	0.0206	0.9927	0.9875	0.9927
Transformer	1	0.0178	0.0305	0.9842	0.9663	0.9842
	2	0.0161	0.0262	0.9881	0.9781	0.9881
	3	0.0162	0.0266	0.9875	0.9751	0.9875
	4	0.0126	0.0243	0.9904	0.9826	0.9904
	5	0.0115	0.0213	0.9922	0.9831	0.9922
TCN	1	0.0122	0.0269	0.9877	0.9806	0.9877
	2	0.0108	0.0217	0.9918	0.9891	0.9918
	3	0.0113	0.0223	0.9913	0.9865	0.9913
	4	0.0115	0.0228	0.9915	0.9888	0.9915
	5	0.0104	0.0201	0.9931	0.9904	0.9931

4.3 Overall comparison of ML models

This section describes the overall performance of the models by averaging all the performance metrics for M1 and M2 input types given in the table 4.

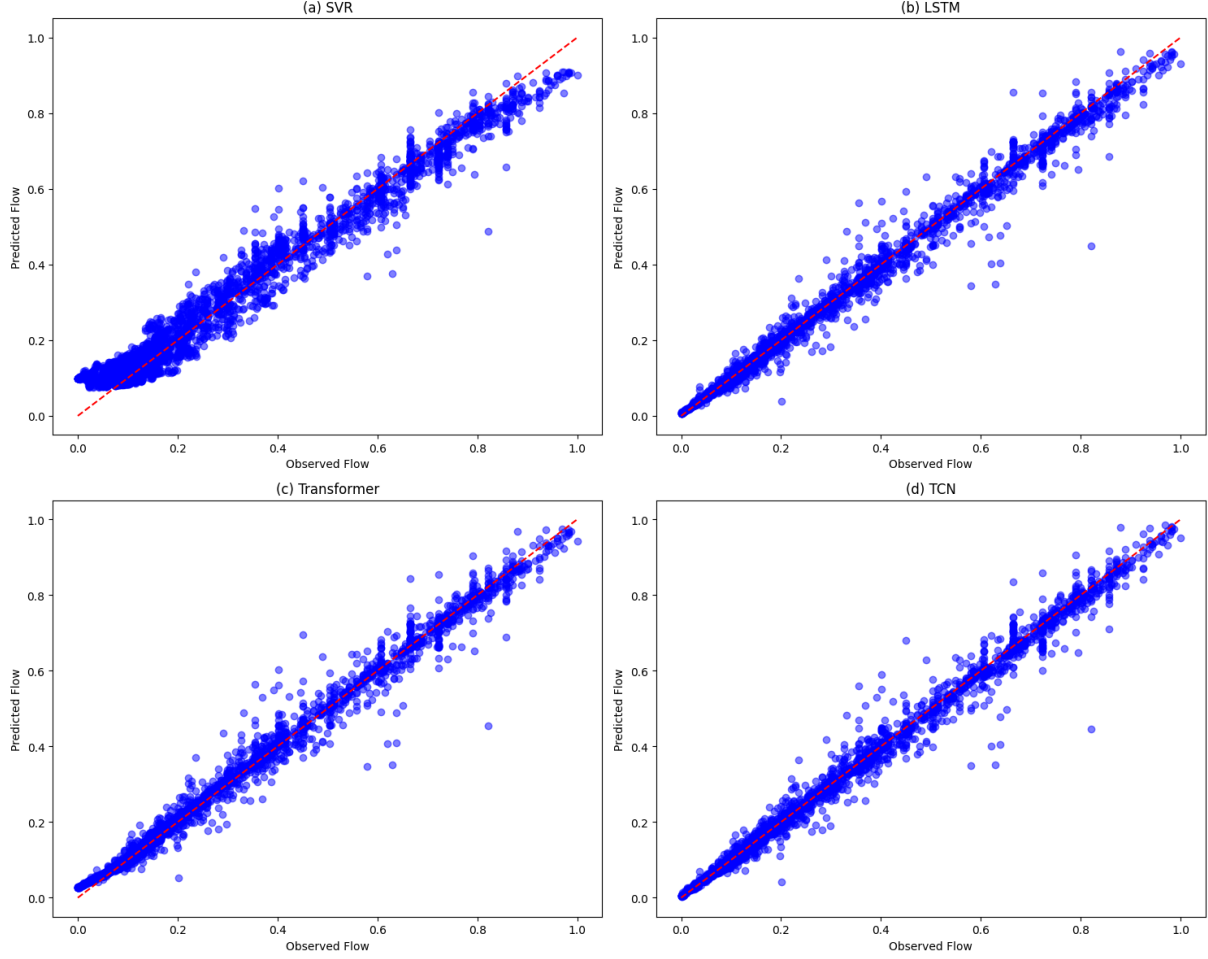


Figure 7: Scatter plot of predicted flow with observed flow (a) SVR model (b) LSTM model (c) Transformer model (d) TCN model. The solid line is the 1:1 line.

Table 4 shows the value of different performance metrics obtained by all the models over five outer folds. The sensitivity analysis of the models for two input conditions, i.e., all four inputs (M1) and all inputs without pressure (M2), is presented in the table. The SVR model performs better for M2 input with average KGE = 93.7% and RMSE = 0.039, as compared to M1 input with an average KGE and RMSE of 91.8% and 0.042, respectively. Similarly, the LSTM model also performs better with M2 input with an average KGE of 98.1% and RMSE of 0.024 and outperforms M1 input with an average KGE and RMSE of 97.5% and 0.025, respectively. Although the transformer model shows better performance for all the other performance metrics in the M2 input except for the average KGE. Specifically, the value of the average KGE with M2 input is 97.7%, and for the M1 input is 98.6%. Finally, the table 4 concludes that the TCN model has better performance as compared to any other models. Unlike other models, TCN performance is better with M1 input with an average KGE of 99.2%, which is 0.5% greater than M2 on average. However, in this overall comparison section the results we have obtained in input M1, and M2 are almost similar. The result obtained in the study by Thapa et al. [2020a] has a very close RMSE value for the M1 and M2 inputs. The final value of RMSE reached in the aforementioned paper is achieved using the hit and trial method which is questionable whether the RMSE could further be reduced, which could increase the model's performance. Hence our research contributes to the fact that the use of a hyper-parameter tuning can significantly reduce the RMSE value for both M1 and M2 as compared to the hit and trial method used in Thapa et al. Furthermore, the research contributes the addition of the P parameter has no significant contribution to snowmelt prediction.

Table 4: Overall performance Comparison with four inputs M1 and three inputs M2 for various models; where performance metrics are: MAE, RMSE, R^2 , KGE, and NSE

Models	INPUTS	MAE	RMSE	R^2	KGE	NSE
SVR	M1	0.032	0.042	0.970	0.918	0.970
	M2	0.030	0.039	0.974	0.937	0.975
LSTM	M1	0.013	0.025	0.989	0.975	0.989
	M2	0.012	0.024	0.990	0.981	0.990
Transformer	M1	0.018	0.028	0.987	0.986	0.987
	M2	0.015	0.026	0.989	0.977	0.989
TCN	M1	0.011	0.023	0.991	0.992	0.991
	M2	0.011	0.023	0.991	0.987	0.991

4.4 Testing time

In this subsection, we have computed the testing time in each outer fold for different models and averaged its testing time as shown in table 5. We have not reported the testing time of the SVR model based on several considerations, including the computational efficiency of SVR relative to other models evaluated in our study. In comparison to more complex models like LSTM, TCN, and Transformers. The training and testing times for SVR are generally lower due to its simpler architecture and optimization algorithms, which are specifically designed for efficiency in high-dimensional feature spaces. By focusing our testing time analysis on LSTM, TCN, and Transformers, we aimed to highlight the computational characteristics of these newer and more sophisticated models, which are of greater interest to the research community given their potential for handling complex sequential data. The table shows the TCN model has a faster computational testing time than any other DL architecture, as shown in table 5.

Table 5: Model testing time (seconds) comparison

Models	Testing Time (seconds)
LSTM	0.694
Transformer	0.652
TCN	0.372

5 Discussion

Snowmelt modeling is useful for estimating the water availability for reservoir management, water supply, irrigation, and hydroelectricity projects, so it is necessary to explore new techniques for the appropriate snowmelt forecasting. Taking Langtang Basin as an example, we applied SOTA DL architectures such as Transformer and TCN for the snowmelt prediction which has not been used for the snowmelt forecasting. In this study, model efficiency (in terms of NSE) for SVR, LSTM, Transformer, and TCN were 97.5%, 99%, 98.9%, and 99.1%, respectively. In a previous study by Uysal et al. [2016], the ANN model was used for the snowmelt runoff prediction in the Upper Euphrates Basin of Turkey and achieved 93% model efficiency which is lower than the ML model used in this study. In the study by Le et al. [2019], the LSTM model has a model efficiency of 99.2% comparable to our LSTM model result used in this study.

Kratzert et al. [2018] used two-layered LSTM whereas Le et al. [2019] and Fan et al. [2020] used single-layer LSTM, these studies did not evaluate the different numbers of LSTM layers. In a study by Thapa et al. [2020a], the performance of the LSTM model for the different hidden layers, optimizers, and window size was evaluated, however, the manual

tuning of hyper-parameter without the use of any systematic approach is questionable for achieving the best performance of the model. In our study, we use a Keras tuner for the hyper-parameter tuning, so that each combination of the hyper-parameter will be evaluated to find the best prediction model. After hyper-parameter tuning, we found that two-layered stacked performs better than single-layer LSTM. In this study, we have compared the performance of five optimizers and found that Adam with a learning rate of 0.004 performs better. The performance of all the models for each fold is shown in Table 2 and 3.

Most of the studies employing ML models (Uysal et al. [2016], Kratzert et al. [2018], Fan et al. [2020], Thapa et al. [2020a]) does not exploit the recent SOTA deep learning architectures. In the past endeavor, traditional approaches such as SVR, ANN, and LSTM were used for snowmelt forecasting. In our study, we have used recent DL architectures such as Transformer and TCN, these models have claimed that they are capable of capturing long-range dependencies and outperformed all other traditional ML models. We found that the TCN model has superior performance than other models with a model efficiency of 99.1% in terms of NSE. In the model development phase input selection is an important task, but it is often neglected. In a study by Thapa et al. [2020a], it was found that gamma testing helps to determine the appropriate input combination, so in this study, we have compared each model on different input types and evaluated that the models with M2 input have better performance which is shown in table 4. However, the TCN model with M1 input has a better model efficiency while all other models are performing better in M2 input which might be due to variations in spatial and temporal resolution or the relevance of the input variable.

Ground truth observation in the Himalayan basins is a difficult task due to the high elevation difference between the stations. Hence, the remotely sensed SCA and meteorological products are important assets. This study proves the applicability of the ML model in snowmelt forecasting using remotely sensed SCA in data-scarce basins. This study signifies that the SOTA models are more applicable for snowmelt forecasting as they capture long-range dependencies and are more generalizable.

6 Conclusion

In this work, we have investigated snowmelt-driven streamflow forecasting using ML techniques. The experiments were conducted over the HKH area. We compared the performance of SVR, LSTM, Transformer, and TCN architectures in M1 and M2 input for forecasting snowmelt. We have found that using M1 and M2 input does not give any significant performance differences also in this research work we have highlighted the importance of the hyperparameter tuning to obtain a better model. The TCN architecture showcased superior performance compared to other ML techniques. Furthermore, we disclosed the benefits of nested CV for performance evaluation on different ML techniques for snowmelt-driven streamflow forecasting. This contribution provides valuable insights for decision-makers in environmental monitoring and resource management, emphasizing the importance of leveraging advanced deep learning architecture for more reliable and precise forecasting.

References

- Prajwal K Panday, Karen E Frey, and Bardan Ghimire. Detection of the timing and duration of snowmelt in the hindu kush-himalaya using quikscat, 2000–2008. *Environmental Research Letters*, 6(2):024007, 2011.
- Philippus Wester, Arabinda Mishra, Aditi Mukherji, and Arun Bhakta Shrestha. *The Hindu Kush Himalaya assessment: mountains, climate change, sustainability and people*. Springer Nature, 2019.
- ASCE Task Committee on Application of Artificial Neural Networks in Hydrology. Artificial neural networks in hydrology. i: Preliminary concepts. *Journal of Hydrologic Engineering*, 5(2):115–123, 2000a.
- ASCE Task Committee on Application of Artificial Neural Networks in Hydrology. Artificial neural networks in hydrology. ii: Hydrologic applications. *Journal of Hydrologic Engineering*, 5(2):124–137, 2000b.
- Mattia Callegari, Paolo Mazzoli, Ludovica De Gregorio, Claudia Notarnicola, Luca Pasolli, Marcello Petitta, and Alberto Pistocchi. Seasonal river discharge forecasting using support vector regression: a case study in the italian alps. *Water*, 7(5):2494–2515, 2015.
- Ludovica De Gregorio, Mattia Callegari, Paolo Mazzoli, Stefano Bagli, Davide Broccoli, Alberto Pistocchi, and Claudia Notarnicola. Operational river discharge forecasting with support vector regression technique applied to alpine catchments: Results, advantages, limits and lesson learned. *Water resources management*, 32:229–242, 2018.
- Gökçen Uysal, Aynur Şensoy, and A Arda Şorman. Improving daily streamflow forecasts in mountainous upper euphrates basin by multi-layer perceptron model with satellite snow products. *Journal of Hydrology*, 543:630–650, 2016.

- Mohammad Najafzadeh, Elahe Sadat Ahmadi-Rad, and Daniel Gebler. Ecological states of watercourses regarding water quality parameters and hydromorphological parameters: deriving empirical equations by machine learning models. *Stochastic Environmental Research and Risk Assessment*, pages 1–24, 2023.
- Mohammad Najafzadeh and Sedigheh Anvari. Long-lead streamflow forecasting using computational intelligence methods while considering uncertainty issue. *Environmental Science and Pollution Research*, 30(35):84474–84490, 2023.
- D Nagesh Kumar, K Srinivasa Raju, and T Sathish. River flow forecasting using recurrent neural networks. *Water resources management*, 18:143–161, 2004.
- Frederik Kratzert, Daniel Klotz, Claire Brenner, Karsten Schulz, and Mathew Herrnegger. Rainfall–runoff modelling using long short-term memory (lstm) networks. *Hydrology and Earth System Sciences*, 22(11):6005–6022, 2018.
- Xuan-Hien Le, Hung Viet Ho, Giha Lee, and Sungho Jung. Application of long short-term memory (lstm) neural network for flood forecasting. *Water*, 11(7):1387, 2019.
- Hongxiang Fan, Mingliang Jiang, Ligang Xu, Hua Zhu, Junxiang Cheng, and Jiahu Jiang. Comparison of long short term memory networks and the hydrological model in runoff simulation. *Water*, 12(1):175, 2020.
- Francesco Granata, Fabio Di Nunno, Mohammad Najafzadeh, and Ibrahim Demir. A stacked machine learning algorithm for multi-step ahead prediction of soil moisture. *Hydrology*, 10(1):1, 2022.
- Mohammad Najafzadeh, Sajad Basirian, and Zhiqiang Li. Vulnerability of the rip current phenomenon in marine environments using machine learning models. *Results in Engineering*, 21:101704, 2024.
- Samit Thapa, Zebin Zhao, Bo Li, Lu Lu, Donglei Fu, Xiaofei Shi, Bo Tang, and Hong Qi. Snowmelt-driven streamflow prediction using machine learning techniques (lstm, narx, gpr, and svr). *Water*, 12(6):1734, 2020a.
- RGI Consortium, Randolph Glacier Inventory, et al. A dataset of global glacier outlines: Version 6.0. *Global Land Ice Measurements from Space: Boulder, CO, USA*, 2017.
- Silvan Ragettli, Francesca Pellicciotti, Walter W Immerzeel, Evan S Miles, Lene Petersen, Martin Heynen, Joseph M Shea, Dorothea Stumm, S Joshi, and A Shrestha. Unraveling the hydrology of a himalayan catchment through integration of high resolution in situ data and remote sensing with an advanced simulation model. *Advances in Water Resources*, 78:94–111, 2015.
- Natsuko Yasutomi, Atsushi Hamada, and Akiyo Yatagai. Development of a long-term daily gridded temperature dataset and its application to rain/snow discrimination of daily precipitation. *Global Environmental Research*, 15(2):165–172, 2011.
- Samit Thapa, Bo Li, Donglei Fu, Xiaofei Shi, Bo Tang, Hong Qi, and Kun Wang. Trend analysis of climatic variables and their relation to snow cover and water availability in the central himalayas: a case study of langtang basin, nepal. *Theoretical and Applied Climatology*, 140:891–903, 2020b.
- Walter W Immerzeel, P Droogers, SM De Jong, and MFP Bierkens. Large-scale monitoring of snow cover and runoff simulation in himalayan river basins using remote sensing. *Remote sensing of Environment*, 113(1):40–49, 2009.
- Emmy E Stigter, Niko Wanders, Tuomo M Saloranta, Joseph M Shea, Marc FP Bierkens, and Walter W Immerzeel. Assimilation of snow cover and snow depth into a snow model to estimate snow water equivalent and snowmelt runoff in a himalayan catchment. *The Cryosphere*, 11(4):1647–1664, 2017.
- Dorothy K Hall, George A Riggs, Vincent V Salomonson, Nicolo E DiGirolamo, and Klaus J Bayr. Modis snow-cover products. *Remote sensing of Environment*, 83(1-2):181–194, 2002.
- Harald Kling, Martin Fuchs, and Maria Paulin. Runoff conditions in the upper danube basin under an ensemble of climate change scenarios. *Journal of hydrology*, 424:264–277, 2012.
- Vladimir Vapnik. *The nature of statistical learning theory*. Springer science & business media, 1999.
- Sepp Hochreiter and Jürgen Schmidhuber. Long short-term memory. *Neural computation*, 9(8):1735–1780, 1997.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. *Advances in neural information processing systems*, 30, 2017.
- Shaojie Bai, J Zico Kolter, and Vladlen Koltun. An empirical evaluation of generic convolutional and recurrent networks for sequence modeling. *arXiv preprint arXiv:1803.01271*, 2018.