

SEPTEMBER 2025

Ph.D. in Civil Engineering

ŞEYDANUR ŞEBÇİOĞLU MUTLU

REPUBLIC OF TÜRKİYE
GAZİANTEP UNIVERSITY
GRADUATE SCHOOL OF NATURAL & APPLIED SCIENCES

INVESTIGATION OF CLIMATE CHANGE IMPACT ON LARGE
DAM PROJECTS EQUIPPED WITH HYDROPOWER

Ph.D. THESIS
IN
CIVIL ENGINEERING

BY
ŞEYDANUR ŞEBÇİOĞLU MUTLU

SEPTEMBER 2025

**INVESTIGATION OF CLIMATE CHANGE IMPACT ON LARGE
DAM PROJECTS EQUIPPED WITH HYDROPOWER**

Ph.D. Thesis

in

**Civil Engineering
Gaziantep University**

Supervisor

Prof. Dr. Aytaç GÜVEN

by

Şeydanur ŞEBİOĞLU MUTLU

September 2025



©2025[Gaziantep University]

I hereby declare that all information in this document has been obtained and presented in accordance with academic rules and ethical conduct. I also declare that, as required by these rules and conduct, I have fully cited and referenced all material and results that are not original to this work.

Şeydanur ŞEBCİOĞLU MUTLU

ABSTRACT

INVESTIGATION OF CLIMATE CHANGE IMPACT ON LARGE DAM PROJECTS EQUIPPED WITH HYDROPOWER

ŞEBİCİOĞLU MUTLU, Seydanur
Ph.D. in Civil Engineering
Supervisor: Prof. Dr. Aytaç GÜVEN
September 2025
104 pages

Hydrological prediction is crucial for managing water resources, and innovations like machine learning present an opportunity to enhance predictive modeling capabilities. The aim of this study is to compare the usage of the new machine learning method, CatBoost, with traditional methods about investigating prediction streamflow in the future under climate change scenarios and determining the effect of predicted streamflows on hydropower energy production. The investigation found that CatBoost was superior to conventional models. After it was proven that the best-performing model is CatBoost, future projections according to the NorESM2-MM scenarios were calculated using this model. Climate projections are based on simulations from the Coupled Model Intercomparison Project Phase 6 model, utilizing shared socioeconomic pathway (SSP) scenarios. The results show that SSP3-7.0 and SSP5-8.5 scenarios indicate an increasing trend between 2015 and 2100, while SSP1-2.6 and SSP2-4.5 expect a balancing tendency. This suggests that climate change has little effect on the measuring station and its basin and that the flow is increasing positively. On the other hand, it is noted that the estimated energy amount in the SSP1-2.6 and SSP2-4.5 scenarios is lower, particularly in the 2015–2039 timeframe. The amount of energy computed under the SSP3-7.0 and SSP5-8.5 scenarios, however, shows a notable increase.

Key Words: CatBoost, Climate change, CMIP6, Hydrological forecasting, Machine Learning, Hydropower

ÖZET

İKLİM DEĞİŞİMİNİN BÜYÜK ÖLÇEKLİ HİDROELEKTRİK SANTRALLERİNE OLAN ETKİSİNİN İNCELENMESİ

ŞEBİOĞLU MUTLU, Şeydanur
Doktora Tezi, İnşaat Mühendisliği
Danışman: Prof. Dr. Aytaç GÜVEN
Eylül 2025
104 sayfa

Hidrolik tahmin, su kaynaklarının yönetimi için çok önemlidir ve makine öğrenimi gibi yenilikler, tahminsel modelleme yeteneklerini geliştirmek için bir fırsat sunar. Bu çalışmanın amacı, iklim değişikliği senaryoları altında gelecekteki akış tahminlerini araştırmak ve tahmin edilen akışların hidroelektrik enerji üretimi üzerindeki etkisini belirlemek konusunda yeni makine öğrenimi yöntemi CatBoost'un kullanımını geleneksel yöntemlerle karşılaştırmaktır. Araştırma, CatBoost'un geleneksel modellere göre üstün olduğunu bulmuştur. En iyi performans gösteren modelin CatBoost olduğu kanıtlandıktan sonra, NorESM2-MM senaryolarına göre gelecekteki projeksiyonlar bu model kullanılarak hesaplanmıştır.. İklim projeksiyonları, paylaşılan sosyoekonomik yol (SSP) senaryolarını kullanan Coupled Model Intercomparison Project Phase 6 modelinden yapılan simülasyonlara dayanmaktadır. Sonuçlar, SSP3-7.0 ve SSP5-8.5 senaryolarının 2015 ile 2100 yılları arasında artan bir eğilim gösterdiğini, SSP1-2.6 ve SSP2-4.5 senaryolarının ise dengeleyici bir eğilim beklediğini göstermektedir. Bu, iklim değişikliğinin ölçüm istasyonu ve havzası üzerinde çok az etkisi olduğunu ve akışın olumlu yönde arttığını göstermektedir. Öte yandan, SSP1-2.6 ve SSP2-4.5 senaryolarında tahmin edilen enerji miktarının özellikle 2015-2039 zaman diliminde daha düşük olduğu belirtilmektedir. Ancak, SSP3-7.0 ve SSP5-8.5 senaryoları altında hesaplanan enerji miktarı dikkate değer bir artış göstermektedir.

Anahtar Kelimeler: CatBoost, İklim Değişimi, CMIP6, Hidrolojik Tahmin, Makine Öğrenimi, Hidro güç



"Dedicated to my family"

ACKNOWLEDGEMENTS

I would like to thank my supervisor. Prof. Dr. Aytaç GÜVEN for his guidance and support throughout the study. I am thankful for his encouragement and motivation.

I would like to express my love and gratitude to my family for their support, always best wishes.



TABLE OF CONTENTS

	Page
ABSTRACT	i
ÖZET.....	ii
DEDICATION.....	iii
ACKNOWLEDGEMENTS.....	iv
TABLE OF CONTENTS.....	v
LIST OF TABLES	viii
LIST OF FIGURES	ix
LIST OF SYMBOLS	xii
LIST OF ABBREVIATIONS	xiii
CHAPTER 1 INTRODUCTION	1
1.1 The Importance of Predicting Future Streamflows Under Climate Change ...	1
1.2 Climate Change Scenarios and Streamflow Forecasting	2
1.2.1 Applications of Climate Change and Current Forecasts.....	2
1.3 Novelty Statement of Study	4
CHAPTER 2 LITERATURE REVIEW	5
2.1 IPCC Working Principle and Scenarios	5
2.1.1 IPCC Scenarios	6
2.2 NorESM2, CMIP6, GCM and Their Interconnections.....	8
2.2.1 Global Climate Models (GCM) and Climate Projections.....	9
2.2.2 NorESM2 and Its Role in Climate Projections	9
2.2.3 The Interaction of NorESM2 and CMIP6.....	10
2.3 CatBoost.....	10
2.3.1 Definition and Development of CatBoost.....	10
2.3.2 Applications of CatBoost.....	11
2.3.3 The Role of CatBoost in Hydrology	11
2.4 Support Vector Machine	14

2.4.1 Definition and Development of Support Vector Machine.....	14
2.4.2 Applications of Support Vector Machine	14
2.4.3 The Role of Support Vector Machine in Hydrology	15
2.5 Gene Expression Programming (GEP).....	17
2.5.1 Definition and Development of GEP	17
2.5.2 Application of GEP.....	18
2.5.3 The Role of GEP in Hydrology	18
2.6 Ridge Regression.....	21
2.6.1 Definition and Development of Ridge Regression	21
2.6.2 Applications of Ridge Regression	21
2.6.3 The Role of Ridge Regression in Hydrology.....	22
CHAPTER 3 METHODOLOGY	25
3.1 The Methodology of CatBoost.....	25
3.2 The Methodology of Support Vector Machine	28
3.3 The Methodology of GEP	30
3.4 The Methodology of Ridge Regression.....	32
3.5 Shapley Additive Explanations (SHAP) and Its Use with Machine Learning.....	32
3.6 Case Study for Imputation Missing Data	33
3.7 Case Study for Comparison Methods and Projections.....	35
3.8 Handling NorESM2 Data	38
3.9 Downscaling using GEP and SVM techniques	40
3.10 Assessment Metrics.....	41
3.10.1 Root Mean Squared Error (RMSE).....	43
3.10.2 Mean Absolute Error (MAE):.....	43
3.10.3 Root Mean Square Error to Standard Deviation Ratio (RSR)	44
3.10.4 Nash-Sutcliffe Model Efficiency (NSE).....	46
3.10.5 Kling-Gupta Efficiency (KGE).....	47
3.11 Mann-Kendall Test.....	48
3.12 Prediction Energy Production	50

CHAPTER 4 RESULT AND DISCUSSION	53
4.1 Results for Imputation for Missing Data.....	53
4.2 Comparison of Methods Results	63
4.3 Future projection of discharge with CatBoost model under SSP scenarios..	78
4.4 Determination of Energy Production with Predicted Discharges	85
4.5 Discussion	86
CHAPTER 5 CONCLUSIONS.....	93
5.1 Imputation of Missing Data.....	93
5.2 Comparison of Conventional Models and CatBoost.....	93
5.3 Comparison of Energy Production.....	94
REFERENCES.....	97
CURRICULUM VITAE.....	104

LIST OF TABLES

	Page
Table 3.1 The optimal hyperparameters for Catboost models	28
Table 3.2 Station discharges (m^3/s) measure results with missing data.....	34
Table 3.3 Station informations and historical discharges (m^3/s).....	36
Table 3.4 List of the NorESM2 variables (predictors) IDs and corresponding variable names.....	39
Table 3.5 Calculated dam heights results under each scenario according to mass curve analysis.....	51
Table 4.1 Comparison of RMSE and MAE results of train and test period before/after imputation	56
Table 4.2 Results of discharges for (m^3/s) for training period (1988-2006)	67
Table 4.3 Results of discharges(m^3/s) for testing period (2007-2014).....	68
Table 4.4 Differences for testing period.....	68
Table 4.5 Comparison of model performance of observed and predicted Q (m^3/s) for testing period between 2007-2014 years.....	69
Table 4.6 Comparison of predicted discharges (m^3/s) between 2015-2039 years under each scenario with CatBoost	79
Table 4.7 Comparison of predicted discharges (m^3/s) between 2040-2069 years under each scenario with CatBoost	80
Table 4.8 Comparison of predicted discharges (m^3/s) between 2070-2100 years under each scenario with CatBoost	80
Table 4.9 Comparison of future projection under different SSPs with observed Q (m^3/s) between 2015-2100 years.	81
Table 4.10 Results of mean discharges and energy prediction	85
Table 4.11 Comparison of mean monthly discgarhe (m^3/s) change in future projections observed by means of (%) percentage.....	86
Table 5.1 Comparison of energy production differences (%) under SSP scenarios based on observed discharges	96

LIST OF FIGURES

	Page
Figure 2.1 Working Principle of IPCC.....	6
Figure 3.1 Structure of CatBoost algorithm	26
Figure 3.2 A schematic diagram for support vector machine (SVM) training and testing process.....	29
Figure 3.3 An illustration of Gene Expression Programming (GEP) algorithm.....	31
Figure 3.4 Location of D20A036 flow gauge station.....	33
Figure 3.5 Flow diagram of the decision process for model performance and future projections.....	35
Figure 3.6 Mean monthly discharge amounts between 1988 and 2014.....	37
Figure 3.7 Location of Gauge Station and Digital Elevation Model (DEM) contours of the watershed.	38
Figure 3.8 An illustration of location of Station and the nearest dam Menzelet.....	52
Figure 4.1 Scatterplot of train period 1979-1994 result before imputation	54
Figure 4.2 Scatterplot of test period 1995-2014 result after imputation	54
Figure 4.3 Scatterplot of train period 1979-2014 result before imputation	55
Figure 4.4 Scatterplot of train period result after imputation.....	55
Figure 4.5 Scatterplot of test period result after imputation	55
Figure 4.6 Correlation coefficient of inputs with discharge before imputation.....	57
Figure 4.7 Correlation coefficients of inputs with discharge after imputation	58
Figure 4.8 Correlation coefficients using simplest method after imputation with monthly mean discharge	59
Figure 4.9 Average Impact Results Before Imputation.....	60
Figure 4.10 Average Impact Results After Imputation	61
Figure 4.11 Impact results with SHAP value before imputation	62
Figure 4.12 Impact results with SHAP value after imputation	63
Figure 4.13 Correlation coefficient of inputs with discharge.....	65
Figure 4.14 SHAP values average impact on model output magnitude.....	66

Figure 4.15 Impact results and direction of SHAP values	67
Figure 4.16 Observed and predicted monthly mean discharge performance for the testing period by all models	69
Figure 4.17 Scatterplot of observed and predicted Q for training period with CatBoost	70
Figure 4.18 Scatterplot of observed and predicted Q for testing period with CatBoost	71
Figure 4.19 Scatterplot of observed and predicted Q for training period with GEP. 71	
Figure 4.20 Scatterplot of observed and predicted Q for testing period with GEP... 72	
Figure 4.21 Scatterplot of observed and predicted Q for training period with Ridge Regression.....	72
Figure 4.22 Scatterplot of observed and predicted Q for testing period with Ridge Regression.....	73
Figure 4.23 Scatterplot of observed and predicted Q for training period with SVM 73	
Figure 4.24 Scatterplot of observed and predicted Q for testing period with SVM . 74	
Figure 4.25 Comparison of the result of observed and predicted discharges for training period in GEP method.....	74
Figure 4.26 Comparison of the result of observed and predicted discharges for test period in GEP method.....	75
Figure 4.27 Comparison of the result of observed and predicted discharges for training period in Ridge Regression method.....	75
Figure 4.28 Comparison of the result of observed and predicted discharges for test period in Ridge Regression method.....	75
Figure 4.29 Comparison of the result of observed and predicted discharges for training period in SVM method	76
Figure 4.30 Comparison of the result of observed and predicted discharges for test period in SVM method	76
Figure 4.31 Comparison of the result of observed and predicted discharges for training period in CatBoost method	76
Figure 4.32 Comparison of the result of observed and predicted discharges for test period in CatBoost method	77
Figure 4.33 Line graph of observed and predicted mean discharge with default model parameters under SSP1-2.6 scenario	81

Figure 4.34 Line graph of observed and predicted mean discharge with default model parameters under SSP2-4.5 scenario	82
Figure 4.35 Line graph of observed and predicted mean discharge with default model parameters under SSP3-7.0 scenario	82
Figure 4.36 Line graph of observed and predicted mean discharge with default model parameters under SSP5-8.5 scenario	83
Figure 4.37 Comparison of the mean discharge between 2015-2039 years projections in CatBoost model under different SSP scenarios.	83
Figure 4.38 Comparison of the mean discharge between 2040-2069 years projections in CatBoost model under different SSP scenarios.	84
Figure 4.39 Comparison of the mean discharge between 2070-2100 years projections in CatBoost model under different SSP scenarios.	84
Figure 4.40 Comparison of energy generation results.	85
Figure 4.41 The projection of annual discharge (Q) values under the NorESM2-MM SSP1-2.6 scenario by CatBoost model and observed data between 1988-2014 in normal scale with a trendline	87
Figure 4.42 The projection of annual discharge (Q) values under the NorESM2-MM SSP2-4.5 scenario by CatBoost model and observed data between 1988-2014 in normal scale with a trendline	87
Figure 4.43 The projection of annual discharge (Q) values under the NorESM2-MM SSP3-7.0 scenario by CatBoost model and observed data between 1988-2014 in normal scale with a trendline	88
Figure 4.44 The projection of annual discharge (Q) values under the NorESM2-MM SSP5-8.5 scenario by CatBoost model and observed data between 1988-2014 in normal scale with a trendline	89

LIST OF SYMBOLS

Q	Discharge
σ	The standard deviation of the observed data.
μ	The average value of the sum of all datas.



LIST OF ABBREVIATIONS

CatBoost	Categorical Boosting
GEP	Gene Expression Programming
GCM	Global Climate Model
IPCC	Intergovernmental Panel on Climate Change
KGE	Kling-Gupta Efficiency
MAE	Mean Absolute Error
MSE	Mean Square Error
NorESM2	Norwegian Earth System Model Version 2
NSE	Nash-Sutcliffe model Efficiency
RMSE	Root Mean Square Error
RSR	Root Mean Square Error to Standard Deviation Ratio
SHAP	Shapley Additive Explanations
SSP	Shared Socioeconomic Pathways
SVM	Support Vector Machine

CHAPTER 1

INTRODUCTION

1.1 The Importance of Predicting Future Streamflows Under Climate Change

Global water resources and ecosystems are being severely affected by climate change. The water cycle and hydrological processes are directly impacted by these changes, which makes planning and managing water resources more difficult. In order to forecast future climate conditions and possible hydrological changes that might occur under them, climate change scenarios are a crucial tool. In particular, streamflow forecasts are crucial for managing water resources and hydrological modeling (Arnell and Gosling 2013). The significance of forecasting future flows under climate change scenarios from a hydrological standpoint has been investigated in this study.

The water cycle is changing significantly as a result of climate change, rising global temperatures, and an increase in carbon dioxide (CO₂) and other greenhouse gases in the atmosphere. Rising temperatures, altered precipitation patterns, increased evaporation rates, and changes in the snow/rain balance are just a few of the ways that climate change is affecting hydrological processes. In particular, the management of water resources is directly impacted by these changes. Water flows (currents) can be directly impacted by changes in precipitation. By increasing evaporation and evaporation losses, higher temperatures can change how water infiltrates the ground and flows at the surface. River flows and water reserves may also be impacted by changes in the seasonal distribution of water resources brought about by the melting of snow and glaciers. For these reasons, future flow predictions under climate change scenarios are of critical importance for the efficient and sustainable management of water resources (Milly et al., 2005).

Flow forecasts are an important tool used to determine the future quantities of water flows in a region. Flow forecasts play a critical role in many areas such as water resources management, dam operation, water supply, hydroelectric energy

production, agricultural irrigation, and flood risk management. With climate change, the accuracy of flow predictions is becoming even more important for water resource management. Hydrological modeling is a fundamental method used in making flow predictions. These models attempt to predict future flows by considering current climate data, land use, watershed characteristics, and other factors. Under climate change scenarios, it is important to accurately integrate the effects of climate changes for these models to provide accurate results. Climate change scenarios allow for the modeling of changes in parameters such as future temperature, precipitation, and evaporation. Thus, flow predictions can be made according to a future scenario that differs from the current climate conditions. For instance, hydrological modeling used to manage a region's water resources can yield crucial data based on climate change scenarios, including predictions about future water reserve declines, the availability of enough water for irrigation, and changes in reservoir levels. Decision-making based on this information facilitates the more effective management of water resources.

1.2 Climate Change Scenarios and Streamflow Forecasting

Climate change scenarios predict future climate conditions based on various climate models. These scenarios vary according to different greenhouse gas emission levels, economic growth scenarios, and environmental factors. Climate change scenarios are generally created using parameters such as Global Warming Potential and Carbon Intensity. These scenarios provide critical data for hydrological models to make future flow predictions. Flow predictions made under climate change scenarios enable the management of future water resources and the development of adaptation strategies. For example, under high greenhouse gas emission scenarios like "RCP 8.5," an increase in temperatures and changes in precipitation patterns are expected. In this case, problems such as increased or decreased flows, exceeding the capacity of dams, or the depletion of water reserves may be encountered. Such scenarios necessitate the implementation of measures for the efficient management of water resources (IPCC, 2014).

1.2.1 Applications of Climate Change and Current Forecasts

There are many application areas for flow predictions made under climate change scenarios. Among these areas are flood risk management, water supply, agricultural

irrigation, hydroelectric energy production, water quality management, and biodiversity conservation strategies.

- **Flood Risk Management**

Climate change can lead to increased or irregular rainfall. This situation can particularly increase the risk of floods caused by sudden and heavy rainfall. Flow predictions allow for the early identification of such flood events and the implementation of necessary precautions. For example, by predicting how water flows will change in the future, settlement areas and infrastructure projects can be planned accordingly.

- **Agricultural Irrigation**

Agricultural production is largely dependent on water resources. Climate change can increase the need for irrigation water or limit water resources. Flow forecasts provide information on how water will be distributed in the future, helping to manage irrigation strategies more effectively. Additionally, irrigation water planning can be carried out during drought periods.

- **Hydroelectric Energy Production**

Hydroelectric power plants generate energy based on water flows. Climate change can affect this energy production by impacting river flows. Flow forecasts allow for predicting the energy production capacity of hydroelectric plants based on future water levels. This way, energy supply can be secured.

- **Water Quality Management**

Water quality is affected by factors such as the temperature, pH level, and pollutant concentration of the water. Flow predictions can help estimate the distribution and concentration of pollutants carried by the water. This allows for the determination of necessary strategies for the protection and improvement of water quality.

Predicting flow in the context of climate change is crucial from a hydrological standpoint. These forecasts offer vital information for better water resource management, ecosystem preservation, and human settlement safety. Flood risk, energy production, and water quality are just a few of the many areas where flow

forecasts are crucial for strategic decision-making. To comprehend the effects of climate change and make precise forecasts, sophisticated hydrological modeling and analyses based on climate change scenarios are therefore required. The significance of forecasting future flows from a hydrological standpoint under climate change scenarios has been investigated in this study.

1.3 Novelty Statement of Study

Numerous studies have already been conducted on the impact of climate change on streamflow, utilising both traditional statistical methods and popular machine learning algorithms like long short-term memory networks, support vector machines, and artificial neural networks. Newer-generation algorithms, especially CatBoost, are still not fully explored for long-term streamflow forecasting under climate change scenarios. The direct consequences for hydropower production have received less attention in the majority of earlier studies, which have also mostly concentrated on discharge projections. This research fills in these gaps by:

- CatBoost's superior predictive performance over traditional and evolutionary algorithms (e.g., SVM, GEP, Ridge Regression) is demonstrated for hydrological modelling under CMIP6-based SSP scenarios.
- Highlighting unexpected streamflow increases under high-emission scenarios, contrary to the widely reported global tendency of discharge reduction, thereby emphasizing the importance of basin-specific responses to climate change.
- Establishing a connection between hydrological forecasts and hydropower generation capability, offering a comprehensive viewpoint that integrates planning for renewable energy and water resources management.

Therefore, by addressing a methodological and thematic gap in the literature, the study offers a novel contribution. It offers fresh perspectives on how climate change might open up locally unique opportunities for water and energy security, in addition to showcasing the potential of sophisticated machine learning algorithms to enhance hydrological forecasts.

CHAPTER 2

LITERATURE REVIEW

2.1 IPCC Working Principle and Scenarios

The Intergovernmental Panel on Climate Change, or IPCC for short, is a global organization concerned with climate change research. The United Nations Environment Programme (UNEP) and the World Meteorological Organization (WMO) founded it in 1988. In order to help policymakers fight global climate change, the IPCC gathers scientific data and analyzes how climate change affects ecosystems and people. These reports help scientists, governments, environmental organizations, and other pertinent stakeholders make science-based decisions.

As authors, contributors, and reviewers, thousands of scientists from around the globe voluntarily support the IPCC's work. The IPCC does not pay any of them. A set of guidelines and procedures serve as the foundation for the IPCC's work. At plenary sessions of government representatives, the Panel makes important decisions. The IPCC's work is supported by a central IPCC Secretariat. There are currently three working groups and a task force within the IPCC. Co-Chairs of that Working Group/Task Force are assisted by Technical Support Units (TSUs), which are hosted and funded by the government of the developed nation. Additionally, a TSU could be set up to assist the IPCC Chair in creating the Synthesis Report for an assessment report. Working Groups I, II, and III address "The Physical Science Basis of Climate Change," "Climate Change Impacts, Adaptation, and Vulnerability," and "Mitigation of Climate Change," respectively. Government representatives are the level at which working groups convene in plenary sessions. Creating and improving a methodology for the computation and reporting of national greenhouse gas emissions and removals is the primary goal of the Task Force on National Greenhouse Gas Inventories. In addition to Task Forces and Working Groups, additional Task Groups and Steering Groups may be formed for a shorter or longer period of time to examine a particular issue or topic. The Task Group on Data and Scenario Support for Impact and Climate Analysis (TGICA) is one such instance.

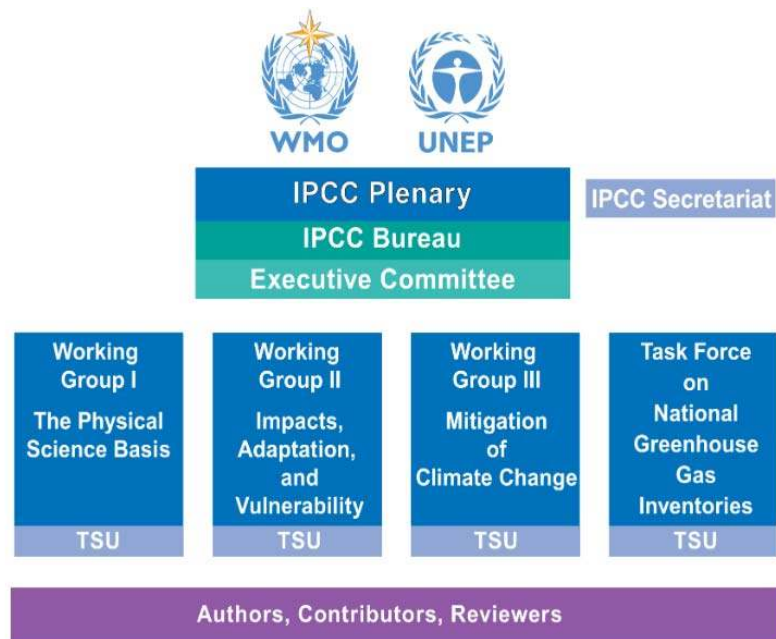


Figure 2.1 Working Principle of IPCC

The working principle of the IPCC is to collect and analyze scientific data on global climate change, make predictions about its impacts, and present these to the public. The panel publishes comprehensive assessment reports every 5-7 years. These reports contain scientific, technical, and socio-economic information related to climate change.

2.1.1 IPCC Scenarios

The IPCC presents a number of future scenarios in its climate change reports. The future changes in atmospheric concentrations of greenhouse gases are modeled by these scenarios. Such projections have been divided into two primary categories since the early 2000s, such as Shared Socioeconomic Pathways (SSP) scenarios and Representative Concentration Pathways (RCP) scenarios (IPCC, 2000).

Representative Concentration Pathways, or RCPs, outline four distinct scenarios for how greenhouse gas emissions will develop in the future. By the year 2100, these pathways predict that the atmospheric concentration of carbon dioxide will reach specific levels. For instance, RCP 2.6 forecasts low levels of greenhouse gas emissions, whereas RCP 8.5 predicts extremely high levels. RCP scenarios are essential for comprehending the magnitude and effects of climate change.

SSP (Shared Socioeconomic Pathways), on the other hand, offers a broader framework. SSP scenarios produce projections for future greenhouse gas emissions by considering global socioeconomic developments, technological advancements, and policy options. SSP scenarios are grouped under 5 main scenarios: low, medium, and high emission scenarios. These scenarios provide projections on how different strategies and societal changes will impact the fight against climate change (IPCC, 2021).

The primary distinction between SSP and RCP is that SSP addresses more general aspects like socioeconomic factors, policy choices, and societal shifts, whereas RCP concentrates exclusively on emission levels. RCP is more concerned with climate physics, whereas SSP offers a more comprehensive approach by accounting for future socio-economic developments. To generate more thorough future projections, SSPs are thus frequently coupled with RCP scenarios (Riahi et al., 2017).

SSPs have been created by the IPCC (Intergovernmental Panel on Climate Change) and are grouped into five main categories: SSP1, SSP2, SSP3, SSP4, and SSP5. Each scenario presents the probabilities of future greenhouse gas emissions by considering the social structure, economic growth, technological developments, environmental policies, and other socio-economic factors.

SSP1 encompasses a population peak and decline (approximately 7 billion by 2100), high income and decreased inequality, efficient land-use regulation, less resource-intensive consumption, including food produced in systems with low greenhouse gas emissions and less food waste, free trade, and eco-friendly technologies and lifestyles. Both mitigation and adaptation challenges are low for SSP1 compared to other pathways (i.e., high adaptive capacity). A low-emission, high-equality society is the goal of SSP1's sustainable world vision. In this case, economic growth takes place within environmental bounds, while global environmental consciousness and sustainable development are at the forefront.

SSP2 accounts for medium income, medium population growth (approximately 9 billion people in 2100), technological advancement, production and consumption patterns that follow historical trends, and only a slow decline in inequality. SSP2 has

medium challenges to adaptation (i.e., medium adaptive capacity) and mitigation compared to other pathways. Current trends are reflected in SSP2, which forecasts modest economic growth accompanied by rising emissions. This scenario typically assumes that policies, technological advancements, and the business world will all continue at a steady pace.

SSP3 encompasses a large population (approximately 13 billion by 2100), low income and persistent inequality, production and consumption of materials, trade barriers, and a slow pace of technological advancement. SSP3 has a low adaptive capacity and significant mitigation and adaptation challenges in comparison to other pathways. SSP3 describes a future characterized by high inequality and regional fragmentation. In this scenario, there is low economic growth, weak environmental policies, and limited technological advancements.

SSP4 accounts for moderate income and population growth (approximately 9 billion people by 2100), but there is a great deal of inequality both within and between regions. In contrast to other pathways, SSP4 presents fewer mitigation challenges but more adaptation challenges (i.e., low adaptive capacity). SSP4 describes a situation where inequalities are increasing globally, with certain regions becoming richer while others are becoming poorer.

SSP5 consists of a population peak and decline (approximately 7 billion in 2100), high income, less inequality, and free trade. Production, consumption, and lifestyles that are resource-intensive are all part of this pathway. In contrast to other pathways, SSP5 presents low adaptation challenges and high mitigation challenges (i.e., high adaptive capacity). SSP5 is a scenario characterized by fossil fuels and rapid economic growth. In this scenario, high energy consumption and increased environmental impacts are expected. (IPCC, 2019)

2.2 NorESM2, CMIP6, GCM and Their Interconnections

Global climate change has become one of the most important environmental and social issues in the world. One of the most powerful tools used to predict the future impacts of climate change is Global Climate Models (GCM). These models simulate the interactions between the atmosphere, oceans, land surfaces, and cryospheric

systems. In this context, NorESM2 stands out as a GCM included in the Coupled Model Intercomparison Project Phase 6 (CMIP6). CMIP6 is an international platform that enables the comparison of global climate projections, and NorESM2 is an important tool that enhances the accuracy of these projections.

2.2.1 Global Climate Models (GCM) and Climate Projections

Global Climate Models (GCMs) are mathematical models that simulate the dynamics and components of the climate. GCMs attempt to predict future climate changes by modeling the interactions of climate components such as the atmosphere, oceans, land surface, and glaciers. GCMs simulate parameters such as future temperature increases, precipitation changes, and sea level rise under different emission scenarios (Eyring et al., 2016). The accuracy of GCMs is related to how closely the simulated climate aligns with the real-world situation. Increasing this accuracy helps climate scientists make better predictions and provides critical data for policymakers (Hurrell et al., 2013).

2.2.2 NorESM2 and Its Role in Climate Projections

Norwegian Earth System Model version 2 (NorESM2) is a GCM developed by the Norwegian Climate Research Institute, simulating numerous climate components and processes. This model simulates the interactions of the atmosphere, oceans, land surfaces, and glacial systems in detail, thereby projecting the impacts of climate change more accurately (Bentsen et al., 2013). NorESM2, within the framework of CMIP6, predicts phenomena such as future temperature increases, precipitation changes, and sea level rise by considering greenhouse gas emissions, water vapor, aerosols, and other climate parameters. These projections provide an important resource for developing strategies to combat climate change. The role of NorESM2 in CMIP6 allows for a more accurate simulation of broad climate changes such as global warming, ocean currents, and biodiversity. NorESM2, as a GCM, is based on similar fundamental principles as other GCMs, but there may be differences in the data and simulations it offers. GCMs generally simulate the interactions between the atmosphere, ocean, and land surfaces, while NorESM2's distinguishing features include the use of more precise physical parameters and the modeling of detailed atmosphere-ocean interactions (Bentsen et al., 2013). The place of NorESM2 in climate projections offers the possibility of comparison with other simulations

conducted with GCMs. Additionally, the simulations of NorESM2 within the CMIP6 framework allow for more comprehensive and reliable results regarding climate change by comparing with the projections of other models.

2.2.3 The Interaction of NorESM2 and CMIP6

CMIP6 is an international research framework that enables the comparison of different climate models and serves as an important data source for climate projections. NorESM2, as one of the models included in the CMIP6 framework, contributes to more accurately predicting global climate changes. CMIP6 tests the accuracy of models by comparing simulations conducted under different emission scenarios and provides important information on their reliability for global climate projections. In this context, the projections provided by NorESM2, when compared with other models offered by CMIP6, help to understand the potential impacts of climate change more comprehensively (IPCC, 2021).

As a result, NorESM2, CMIP6, and GCMs are important complementary tools in understanding climate change projections. GCMs, by modeling the dynamics of the climate, simulate future temperature increases and other climate changes, while NorESM2 stands out as an advanced model that enhances the accuracy of these projections. CMIP6 plays an important role in testing the reliability of these projections by comparing different climate models. The role of NorESM2 within the CMIP6 framework provides a critical resource for understanding global climate change and developing effective strategies to address it.

2.3 CatBoost

2.3.1 Definition and Development of CatBoost

CatBoost is a machine learning algorithm developed by Yandex, aimed at achieving effective results, especially in large and complex datasets. It was first announced in 2017 and has rapidly gained popularity since then. It takes its name from the term "Categorical Boosting" because the algorithm is particularly known for its superiority in processing categorical data. Yandex decided to develop this algorithm to overcome the performance issues that arise with existing boosting algorithms, especially when working with categorical data (Prokhorenkova et al., 2018). The foundation of this

algorithm is based on gradient boosting methods using decision trees. However, unlike older boosting algorithms, CatBoost has the ability to directly handle categorical variables. This feature eliminates the need for traditional algorithms to convert categorical data, allowing for more accurate and faster predictions. The development of CatBoost has been continuously improved by Yandex's research team and has, over time, reached a wide user base in both academic and industrial fields worldwide (Prokhorenkova et al., 2018).

2.3.2 Applications of CatBoost

CatBoost is an algorithm used in a wide range of applications. Especially in fields such as finance, healthcare, retail, and natural language processing, remarkable results have been achieved. In the finance sector, the effectiveness of CatBoost has been proven in applications such as credit risk analysis, fraud detection, and algorithmic trading. CatBoost, especially when working with large datasets, offers high accuracy rates and fast model training processes (Prokhorenkova et al., 2018). In addition, it is widely used in applications such as disease prediction and genetic data analysis in the healthcare sector. In this field, CatBoost's high-accuracy predictions provide significant advantages in terms of early diagnosis and treatment planning. CatBoost is also effectively used in commercial applications such as customer segmentation and sales forecasting in the retail sector. However, the high parallelism support and acceleration techniques offered by the algorithm ensure that operations performed on large datasets are carried out efficiently and quickly.

2.3.3 The Role of CatBoost in Hydrology

In important areas such as hydrology, water resources management, and climate change, the use of machine learning methods to make accurate predictions is becoming increasingly widespread. In this context, CatBoost emerges as a powerful tool in applications such as hydrological modeling and water resources management. Especially since hydrological data often exhibit complex, large, and heterogeneous structures, predictions made using traditional methods may be inadequate. At this point, the advantages offered by CatBoost become evident. CatBoost allows for high-accuracy predictions in the analysis of hydrological data such as water levels, rainfall amounts, and flow rates, thanks to its ability to directly process categorical data and

the effective use of gradient boosting methods. Additionally, CatBoost's overfitting prevention features ensure that long-term hydrological predictions are more robust and reliable. For example, CatBoost offers an effective model that can be used in water resources management based on climate change scenarios, flood predictions, and water quality analyses (Prokhorenkova et al., 2018). These features enable more accurate and reliable results in hydrological modeling processes, allowing for more informed decisions in water resources management. CatBoost is a gradient boosting algorithm that performs particularly well in processing categorical data. In the field of hydrology, the use of this algorithm is still limited, but some studies are progressing in this direction.

- **Streamflow Forecasting**

For water resource management in mountainous catchments, precise daily streamflow forecasting is essential. Szczepanek (2022) evaluated CatBoost against other machine learning models, such as XGBoost and LightGBM, in order to predict daily streamflows in the Polish Skawa River catchment. When trained using precipitation and lagged streamflow data, CatBoost performed better than other models in terms of prediction accuracy, according to the study. Accurate forecasts depended on streamflow data from upstream cross-sections, according to the feature importance analysis.

- **River Network Classification**

River network classification, a task crucial to hydrological modeling and cartography, has also been tackled with CatBoost. In a study by Wang and Qian (2023), river networks were automatically classified by simulating expert classification processes using CatBoost. To improve interpretability, the model used the Shapley additive explanation (SHAP) method and integrated topological, geometric, and semantic aspects of river networks. This method showed how well CatBoost handled intricate hydrological data structures.

- **Flood Prediction**

Another area where CatBoost has shown promise is in flood prediction. Using rainfall data, Kundra et al. (2023) created a forecasting system and trained it using a

variety of machine learning methods, such as CatBoost. According to the study, the CatBoost model outperformed other models such as Multi-Layer Perceptron (MLP) and Extra-Tree classifiers, achieving an accuracy of 98.34% in forecasting flood occurrences. The potential of CatBoost in early flood warning systems is highlighted by its high accuracy.

In the Sakarya Basin, which has a semi-arid climate in Turkey, Kilinc et al. (2023) used a hybrid machine learning model created with CatBoost and Genetic Algorithm for daily river flow prediction. The outcomes demonstrate that CatBoost performed better than alternative techniques like LSTM and Linear Regression.

The CatBoost algorithm was assessed by Ogundolie et al. (2024) for flood prediction in the Osun River Basin. Metrics like accuracy, precision, and recall were used to test the model, and it achieved a 92.48% accuracy rate.

- **Streamflow Prediction in River Basins**

Mishra et al. (2024) used a Genetic Algorithm and CatBoost to forecast streamflow in the Lower Godavari River basin. Streamflow patterns were successfully captured by the CatBoost-Genetic Algorithm combination, indicating the algorithm's resilience in hydrological time-series modeling, according to the study.

- **Other studies using Catboost in Turkey**

At the XII. National Hydrology Congress (2024), a study titled "Performance Analysis of Global Non-Hydrostatic Numerical Weather Prediction Models and CatBoost Components Across Turkey" was presented. In this study, the accuracy of hydrological predictions was examined through the combination of numerical weather prediction models and the CatBoost algorithm.

Additionally, a review published by Eker et al. (2023) provided a general evaluation of the use of machine learning algorithms in studies on forest soil and hydrology. In this study, it was noted that there are a limited number of works in the literature on the use of machine learning techniques in the category of forest soil and hydrology, and that these studies generally address the topics of soil moisture mapping and modeling.

In hydrological prediction tasks, such as streamflow forecasting, river network classification, flood prediction, and rainfall forecasting, the reviewed studies show that CatBoost is an effective tool. Hydrological modeling and water resource management benefit greatly from its high accuracy and capacity to handle complex data structures.

2.4 Support Vector Machine

2.4.1 Definition and Development of Support Vector Machine

Support Vector Machine (SVM) was developed in the mid-1990s by statistician and machine learning expert Vapnik and Cortes (Cortes and Vapnik, 1995). SVM is primarily a classification algorithm, but over time it has also found extensive use in regression and other machine learning problems. SVM has achieved great success, especially when working with non-linear data, by using kernel methods that allow the data to be separated linearly (Vapnik, 2013). The main objective of SVM is to find the decision boundary (hyperplane) that provides the widest marginal gap (margin) to separate the data into two classes. The developers' goal was to achieve high-accuracy classifications by clearly separating the boundaries between the classes. SVM has succeeded in defining non-linear boundaries by transforming data points into a high-dimensional feature space. This technique has proven to be particularly successful in datasets where classes cannot be linearly separated and contain complex structures. Since the late 1990s, SVM has been widely used in both industrial and academic fields due to its effective results on large datasets.

2.4.2 Applications of Support Vector Machine

Support Vector Machine is a powerful machine learning algorithm used in a wide range of applications. Its most common use is in classification problems; however, it is also effectively used in different areas such as regression, anomaly detection, and dimensionality reduction. Especially, significant successes have been achieved in areas such as text classification, face recognition, biological data analysis, financial market predictions, and image recognition (Cortes and Vapnik, 1995). The high accuracy of SVM allows it to perform effective classifications even on large datasets. Similarly, in facial recognition and biometric security systems, SVM provides high-accuracy results along with feature extraction. SVM is also effectively used in fields

with complex and nonlinear data, such as financial data analysis. The success of SVM can be further enhanced by the correct selection of features and the use of kernel functions. In addition, as the size and complexity of the dataset increase, the classification capacity offered by SVM becomes even more pronounced. In conclusion, SVM is a machine learning algorithm that demonstrates superior performance by effectively modeling non-linear and complex data structures in many different fields.

2.4.3 The Role of Support Vector Machine in Hydrology

In critical areas such as hydrology, water resources management, and climate change, the need for machine learning techniques to make accurate predictions is increasing every day. In this context, Support Vector Machine (SVM) emerges as an important tool in hydrological predictions. SVM can be adapted to hydrological models, especially due to its ability to work with large, complex, and nonlinear datasets. Hydrological datasets encounter nonlinear relationships and high-dimensional data because they include measurements of precipitation, water level, flow, and other environmental factors. SVM has the ability to make high-accuracy predictions when modeling such complex structures. SVM is an effective tool, especially in water quality predictions, flood predictions, and water resources management, for defining nonlinear relationships. Thanks to the powerful kernel functions of this algorithm, even nonlinear relationships in hydrological predictions can be accurately modeled. Additionally, SVM offers effective regularization techniques to minimize the risk of overfitting, which ensures more reliable results in long-term predictions. Allows SVM to make more accurate and reliable predictions in hydrological modeling processes, enabling more informed and efficient decision-making in water resources management.

Support Vector Machine (SVM) is a powerful machine learning algorithm widely used in classification and regression analysis. In the field of hydrology, the use of SVM is being researched, particularly in areas such as water demand forecasting and precipitation prediction.

- **Estimation of Water Demand**

Accurately predicting water demand is critically important for the management and planning of water resources. In a study conducted by Msiza et al. (2007), the effectiveness of artificial neural networks (ANN) and SVM in water demand time series forecasting was compared. The results indicated that both methods were successful in predicting water demand, but ANN generally performed better.

- **Rainfall Forecast**

Rainfall forecasting is an important component in hydrological modeling. Hussein et al., (2020) made regional rainfall predictions by classifying large-scale rainfall maps using SVM. This approach has compared the rainfall amounts in different regions thanks to SVM's ability to work with large datasets.

- **Streamflow Prediction**

Essam et al., (2022) research focused on developing a universal model for streamflow prediction across rivers in Peninsular Malaysia. The study employed SVMs to predict streamflow, highlighting the model's adaptability to different river systems within the region.

Support Vector Machine (SVM) was used in Sharma and Goel's (2024) study to forecast the Tehri Dam's daily and ten-day streamflow. The model's potential for streamflow forecasting in the area was demonstrated by the evaluation of its performance using a variety of statistical metrics.

In the South Tyrol region of Italy, Dalla Torre et al. (2024) created a data-driven model based on SVM to forecast short-term flow. In intricate alpine areas, this model improves the precision of flow forecasts.

- **Groundwater Level Prediction**

Radhakrishnan and CA (2023) developed models combining M5 Model Tree and SVM to simulate groundwater levels. The research aimed to enhance prediction accuracy by integrating these models, offering insights into groundwater behavior in the study area.

- **Flood Simulation**

Langhammer, J. (2023) coupled SVM models with a hydrometeorological wireless sensor network to simulate various flood events in a montane basin. The approach demonstrated the feasibility of using SVMs in conjunction with sensor networks for real-time flood simulation.

- **Forecast Drought and Water Scarcity**

SVM was used by Coşkun and Akbaş (2024) to forecast drought and urban water scarcity in Bursa, Turkey. The study assesses how population growth and climate change affect water resources.

Support vector machines have shown promise in hydrology for a number of prediction tasks, such as rainfall prediction, flood simulation, groundwater level prediction, and streamflow forecasting. The 2020–2025 reviewed studies demonstrate how flexible and reliable SVMs are when processing intricate hydrological data. Future studies might concentrate on combining SVMs with real-time data sources and additional machine learning methods to improve hydrology's predictive power.

2.5 Gene Expression Programming (GEP)

2.5.1 Definition and Development of GEP

Gene Expression Programming (GEP) is a type of evolutionary algorithm, developed in 2001 by Cândida Ferreira, which is a combination of genetic algorithms and evolutionary programming (Ferreira, 2001). GEP, based on genetic algorithms, develops computer programs and mathematical models through evolutionary paths, and this process works similarly to the selection, recombination, and mutation processes in biological evolution. Traditional genetic algorithms generate solutions at the genetic code level, while GEP takes a step further to develop solutions in a more complex and effective manner, using a genetic structure to express these sets of solutions. The most important feature of GEP is the representation of chromosomes in a form that carries genetic information and directly produces computer programs. Initially limited to evolutionary algorithms, this model has found applications in many areas over time, such as regression analysis, function discovery,

and data mining. Ferreira (2001), in his studies developing GEP, aimed to create more effective programs using genetic structures and applied this methodology to different problem areas. This technology is similar to genetic algorithms but offers more flexibility and depth, playing a significant role in complex problem-solving processes.

2.5.2 Application of GEP

Gene Expression Programming is used in various engineering, scientific, and commercial fields. These fields are particularly complex analyses that require mathematical modeling, function discovery, data mining, and optimization problems. One of the areas where GEP is particularly preferred is the modeling of complex relationships and problems that need to be expressed as nonlinear functions. In fields such as natural language processing, financial market analysis, bioinformatics, and environmental modeling, GEP offers effective results. One of the strengths of GEP in genetic programming is its ability to represent functions using genetic codes in evolutionary processes aimed at solving very complex problems. For example, in the field of biotechnology, GEP is used to model protein structures and make predictions related to genetic sequences. Additionally, in engineering design, it finds extensive application in determining material properties and solving optimization problems. In the finance sector, GEP is used as an effective tool for tasks such as stock price predictions and risk analysis. In conclusion, GEP is a powerful tool for modeling complex, nonlinear, and dynamic systems across a wide range of applications.

2.5.3 The Role of GEP in Hydrology

In critical areas such as hydrology, water resources management, and environmental modeling, making accurate and reliable predictions is of great importance for sustainable water resources management. Gene Expression Programming (GEP) is being increasingly used in hydrological modeling and forecasting. Hydrological systems often involve complex and nonlinear structures, which can make accurate predictions difficult using traditional methods. GEP is used to model nonlinear relationships in hydrological predictions and to effectively explain the dynamic characteristics of the system. For example, the relationships between hydrological variables such as precipitation, water level, and flow can be accurately discovered

using GEP. The evolutionary structure of GEP allows it to extract meaningful relationships from data sets, making it applicable in water resources management, flood predictions, and water quality analyses. Additionally, GEP can offer high accuracy rates with fewer parameters compared to other methods in hydrological modeling. However, the flexible structure of GEP allows it to adapt to various hydrological conditions and work effectively with different types of data. These features make GEP a powerful tool in hydrological forecasting and enable water resources management to be carried out more efficiently and reliably. These advantages in the field of hydrological modeling lead to the increasing adoption and use of GEP (Ferreira, 2001). Genetic Programming (GP), as a member of the family of evolutionary algorithms, is used in various applications in the field of hydrology. GP stands out as an effective method, especially in modeling complex and nonlinear relationships.

- **Flood Displacement Modeling**

Önen and Oral (2019) used GEP in modeling flood routing processes. In this study, GP-based models were developed to understand how floods move in riverbeds and how they change over time. These models have helped in predicting and managing floods more accurately.

- **Water Network Hydraulic Calibration**

Gözütok and Özdemir (2004) have also used GEP in the hydraulic calibration of water networks. In this application, GP-based methods have been developed to more accurately determine parameters such as pipe roughnesses and node consumption values in water networks. In this way, the performance of the water networks has been improved and water losses have been reduced.

- **Streamflow Forecasting**

Esha and Imteaz (2020) used GEP to create an Artificial Intelligence (AI) model that forecasts streamflow over the long term in Australia. The model forecasted seasonal streamflow using large-scale climate drivers, showcasing GEP's ability to handle complex climatic data for hydrological forecasts.

Yaman and Önen (2023) used GEP and ANFIS techniques to model the precipitation-runoff relationship in the Berta Suyu Sub-Basin using daily precipitation data gathered from satellite data and observed runoff data. According to the results, both approaches produce predictions with a high degree of accuracy.

In six Pakistani cities with varying climates, Raza et al. (2024) used meteorological data, including minimum, maximum, and average air temperature, average relative humidity, wind speed, and sunshine duration, to predict reference evapotranspiration (RET) using the GEP method.

Gene Expression Programming (GEP) and Linear Genetic Programming (LGP) were used by Güven et al. (2024) to forecast flow and sediment data in a sub-basin of the Euphrates River in Turkey. Data on flow and sedimentation under future climate scenarios (RCP2.6, RCP4.5, and RCP8.5) have been projected by GEP.

- **Water Level Fluctuations**

In order to forecast changes in Urmieh Lake's water level, Karimi et al. (2012) used GEP in conjunction with the Adaptive Neuro-Fuzzy Inference System (ANFIS). The study demonstrated GEP's potential for managing water resources and its efficacy in simulating intricate hydrological systems.

- **Discharge Capacity Estimation**

A GEP-based model was presented in a study by Bonakdari et al. (2021) to forecast the discharge coefficient of triangular labyrinth weirs. The model illustrated the potential of GEP in hydraulic engineering applications by taking into account parameters like vertex angle, Froude number, and crest height ratio.

In hydrological prediction, gene expression programming has shown itself to be a useful tool, providing flexibility and resilience when simulating intricate systems. Applications of GEP in streamflow forecasting, water level fluctuation prediction, discharge capacity estimation, and modeling sustainable building materials are highlighted in the 2020–2025 studies that were reviewed. To improve hydrology's predictive power, future studies might concentrate on combining GEP with additional machine learning methods and real-time data sources.

2.6 Ridge Regression

2.6.1 Definition and Development of Ridge Regression

Ridge regression is a statistical method developed to solve the multicollinearity problem encountered in linear regression models. Multicollinearity refers to a high linear relationship between independent variables, and this situation can make it difficult to estimate the parameters of the regression model, and even compromise the model's stability. Ridge regression addresses this issue by adding a penalty term, making the model healthier and more reliable.

Ridge regression adds an "L2 norm" to the basic formula of linear regression. This norm tries to minimize the sum of the squares of the regression coefficients. As a result, the model's coefficients take on small values, thereby reducing the model's complexity. This situation prevents the problem of overfitting and enhances the model's generalization capability. The main advantage of Ridge regression is that, especially in cases with a large number of variables, the model can make more stable and reliable predictions. However, Ridge regression reduces the effects of independent variables rather than completely eliminating them, which allows all variables to remain in the model.

It is believed that ridge regression was first developed in the 1970s by statisticians Arthur E. Hoerl and Robert W. Kennard. It was first introduced to control the magnitude of regression coefficients and to prevent the model from overfitting (Hoerl and Kennard, 1970). The aim of this method is to obtain more accurate and stable predictions, especially when working with data that has multiple linear relationships. Ridge regression is widely used today in the fields of machine learning and statistical modeling. Ridge regression is often preferred in data analysis-based fields such as medicine, finance, and engineering when working with complex datasets.

2.6.2 Applications of Ridge Regression

Ridge regression is an effective solution, especially in cases where multiple linear regression is common and there is high correlation in the datasets. It is widely used in many scientific, engineering, and commercial fields. Some of these fields include financial modeling, bioinformatics, economics, environmental engineering, and social sciences. Especially in cases where multiple independent variables are highly correlated with each other, classical linear regression models can produce high-

variance predictions, which reduces the model's generalizability and accuracy. Ridge regression increases the model's generalization capacity by penalizing parameters in such data. For example, in financial forecasts, asset price predictions, economic indicator analysis, and genetic research, relationships among multiple independent variables are frequently observed. The application of Ridge regression in these areas allows the model to improve its accuracy without overfitting. Ridge regression is also used in the analysis of large datasets such as text mining and natural language processing, where it is common for independent variables (features) to be high-dimensional and have strong relationships with each other.

2.6.3 The Role of Ridge Regression in Hydrology

Hydrology is of critical importance for water resources management and modeling environmental changes. Studies conducted in this field require accurate modeling of data such as water level, flow, and precipitation. Ridge regression is an important tool in hydrological predictions because hydrological datasets often contain multicollinearity. In the data used to understand and predict the water cycle, there may be high correlations between different environmental factors. For example, the relationship between rainfall amount and water level or the relationship between temperature and humidity. In such highly correlated datasets, classical linear regression models can carry a high variance and overfitting risk. Ridge regression reduces this risk by penalizing the coefficients and makes hydrological models more robust and generalizable. This method has achieved great success in areas such as flood predictions, water level modeling, and flow predictions. For example, in predictions related to water quality, it can be observed that different environmental parameters (such as pH, oxygen level, and temperature) are interrelated. In this case, Ridge regression is a highly effective method for making accurate predictions. Additionally, the fact that Ridge regression provides more reliable results even with smaller datasets enhances accuracy and reliability in hydrological research.

Ridge Regression (L2 regularization) is a regression technique used particularly to address issues of multicollinearity and overfitting. In the field of hydrology, this method has been used in various applications.

Singh et al., (2023) created a statistical downscaling model to forecast daily rainfall from large-scale climate variables by combining Principal Component Analysis, Volterra series, and Ridge Regression. Outperforming conventional downscaling techniques, the model successfully captured non-linear interactions.

Zandi et al. (2023) created a framework for the spatial interpolation of monthly precipitation data in an orographically complex region using local weighted linear Ridge Regression (LWRR). By tackling non-linear relationships and multicollinearity, the model has produced predictions with a higher accuracy than conventional techniques. By tackling non-linear relationships and multicollinearity, the model has produced predictions with a higher degree of accuracy than conventional techniques.

A feature-engineered, time-series, cross-validated Ridge Regression model was introduced in Jisha (2024)'s study to increase the accuracy of weather predictions. Advanced feature engineering techniques were used in this model to improve prediction performance.

Using the Guandi Hydroelectric Power Plant as an example, Chen et al. (2024) integrated sub-models like multiple linear regression (MLR), backpropagation neural network (BPNN), wavelet neural network (WNN), and support vector regression (SVR) with Ridge Regression (RR) to forecast medium- and long-term flows. The results obtained demonstrate that, in comparison to other combination models, the RR-3 model offers higher accuracy.

In order to predict streamflow based on daily discharge and precipitation data, Samadi-Koucheksaraee and Chu (2024) created a hybrid machine learning model called WKELM-R, which integrates Ridge Regression with additional algorithms. Compared to conventional methods, this approach showed better predictive performance.

Yi et al., (2025) used deep learning and machine learning models to forecast flow duration curves in ungauged basins. The research highlighted the effectiveness of these models in overcoming challenges associated with data scarcity.

Ridge regression provides robustness against overfitting and multicollinearity, making it a useful tool in hydrological prediction. Its applications in streamflow forecasting, rainfall downscaling, weather prediction, and flow duration curve estimation are highlighted in the 2020–2025 studies that were reviewed. To improve hydrology's predictive power, future studies might concentrate on combining Ridge Regression with additional machine learning methods and real-time data sources.



CHAPTER 3

METHODOLOGY

3.1 The Methodology of CatBoost

CatBoost is a derivative of gradient boosting algorithms and is fundamentally based on decision trees. However, unlike other boosting methods, CatBoost has a special structure to work more efficiently with categorical data. This algorithm uses a special embedding algorithm to represent each categorical variable. Categorical data can be processed directly without being converted into numerical data. This process helps preserve information that could be lost during the conversion of categorical data. The operation of CatBoost consists of the following basic steps:

Tree-Based Modeling: CatBoost builds a decision tree in each iteration. Each new tree tries to correct the errors learned from the previous trees. This uses the gradient boosting method to increase the model's accuracy.

Ranking and Placement: In the processing of categorical data, CatBoost applies specific ranking techniques for each category. In this way, categorical data is correctly encoded, making it usable during the model's training process (Prokhorenkova et al., 2018).

Regularization Techniques: CatBoost uses various regularization techniques to prevent the model from overfitting. This increases the model's generalization ability, resulting in more reliable outcomes.

Acceleration and Parallel Processing: CatBoost supports parallel processing during the training process. This allows the model to be trained more quickly on large datasets. Additionally, the algorithm is optimized to ensure that the processor cores are utilized at full capacity.

As a result, CatBoost offers a powerful learning algorithm through the combination of decision trees and gradient boosting methods. The efficiency in processing categorical data increases the model's accuracy, while the acceleration techniques in the training process allow it to work efficiently on large datasets. These features make CatBoost a tool capable of making highly accurate predictions, especially in complex datasets.

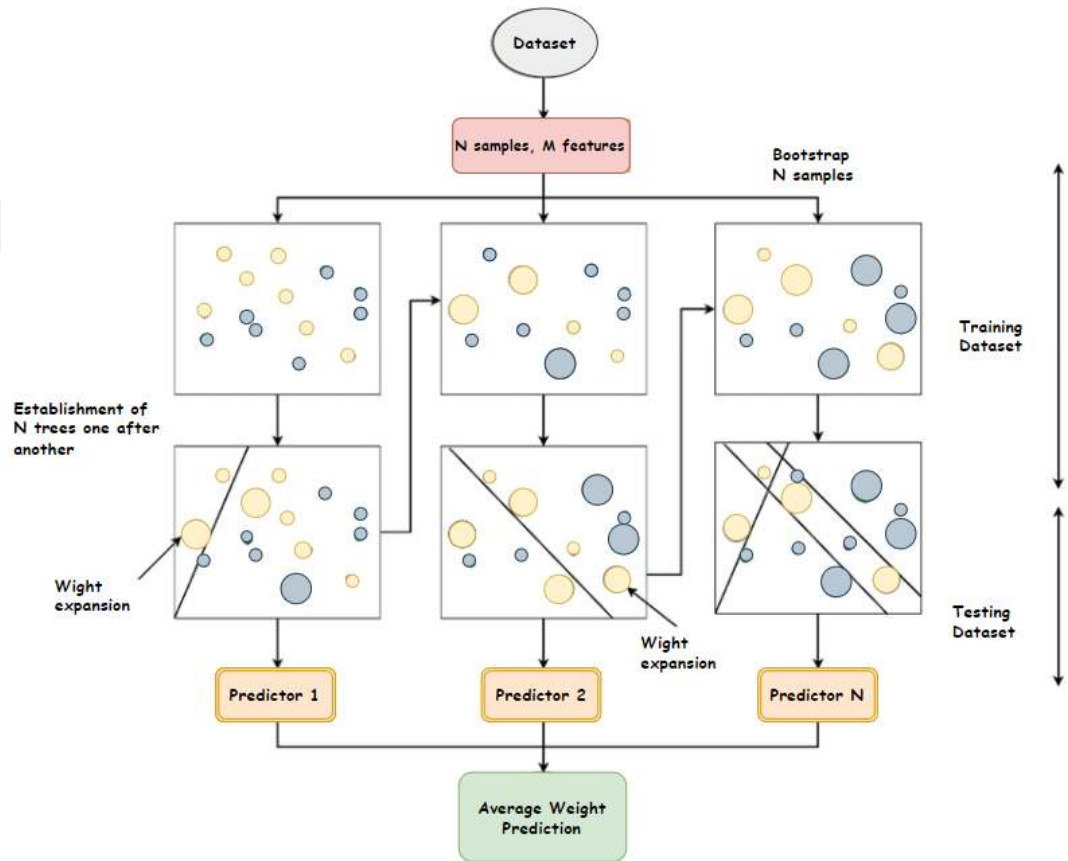


Figure 3.1 Structure of CatBoost algorithm

The Catboost algorithm applies gradient boosting to a decision tree in a sequential manner, using the decision tree as the underlying weak learner. In order to decrease overfitting and enhance the Catboost algorithm's performance, gradient learning data is arranged irregularly. During learning, the Catboost algorithm decreases the amount of movement that is predicted. It avoids overfitting, effectively handles gradient bias and prediction shift, and improves computation accuracy and generalizability. Its efficacy, precision, and capacity to manage categorical features make

CatBoostRegressor noteworthy. Here's how this approach varies from standard gradient boosting decision tree (GBDT) algorithms:

- (1) Training with categorical features instead of preprocessing. The complete dataset is used for CatBoost's training. Target statistics is a useful tactic for managing categorical features while reducing information loss, according to research by Prokhorenkova et al. (2018). In order to determine the average label value for the example with the same category value placed before the supplied one, CatBoost randomly permutes the dataset for every example.
- (2) Features in combination. A new category could be created by combining the features of the existing ones. CatBoost creates new splits for trees in a greedy manner. No combination is looked at for the tree's initial split. CatBoost, on the other hand, uses every predefined combination with every categorical variable in the dataset for the second and subsequent divides. The tree's splits are combined and handled as two-value categories.
- (3) Boosting without bias with categorical features. A deviation from the initial distribution occurs when categorical features are converted to numerical values using the TS approach. With older GBDT techniques, this is a frequent problem.
- (4) Quick Scorer. CatBoost uses the same dividing criterion for all levels and uses oblivious trees as baseline predictors. The size of these trees is appropriate, and they don't overfit. The binary vectors that encode each leaf index in oblivious trees are the same length as the tree depth. This method is frequently used in CatBoost model evaluators to compute model predictions because all binaries make use of statistics, one-hot features, and float (Huang et al., 2019).

In CatBoost, hyperparameters are essential to model performance. How well CatBoost handles categorical features and prevents underfitting or overfitting depends on hyperparameters like learning rate, tree depth, and iterations. The hyperparameters of Catboost models used in this study are given Table 3.1.

Table 3.1 The optimal hyperparameters for Catboost models

Method	Hyperparameters	Selected value
CatBoost	Colsample-by level	0.3
	Depth	2
	Iterations	70
	Bagging	
	Temperature	100
	Border Count	255
	Grow Policy	Symmetric Tree
	Leaf reg	4
	Learning rate	0.11
	Loss Function	Poisson

3.2 The Methodology of Support Vector Machine

Support Vector Machine (SVM) is fundamentally a classification algorithm, but it uses the kernel method to work with non-linear data and to be applicable in a wider range of applications. SVM, while trying to find the decision boundary that will separate the data into two classes, ensures that this boundary is positioned to leave the widest possible gap (margin) between the two classes. The operation of the algorithm consists of the following steps:

Data Transformation and Kernel Functions: SVM initially transforms non-linearly separable data into a high-dimensional space. This transformation is done using kernel functions. Among the most commonly used kernel functions are linear, polynomial, and radial-based functions. These kernel functions allow the data to be separated by non-linear boundaries. Thanks to kernel functions, SVM can work with more complex structures (Cortes and Vapnik, 1995).

Hyperplane and Margin: The main goal of SVM is to find the decision boundary that separates the data into two classes. This boundary is defined as the hyperplane that leaves the widest margin between the data points. The margin is a linear line where the points closest to the data of both classes (support vectors) are located. SVM aims to achieve the most accurate classification by maximizing the size of this margin.

Support Vectors: The success of SVM relies on the data points closest to the decision boundary, known as "support vectors." These points are the most important data points that affect the model's decision-making process. Other data points do not affect the model's decision, which allows SVM to operate more efficiently.

Regularization: SVM applies regularization techniques to minimize the overfitting problem. This regularization is achieved with the C parameter, which determines how close the model will be to linear classification. A low C value provides a wider margin, while a high C value results in tighter classifications with a smaller margin (Vapnik, 2013).

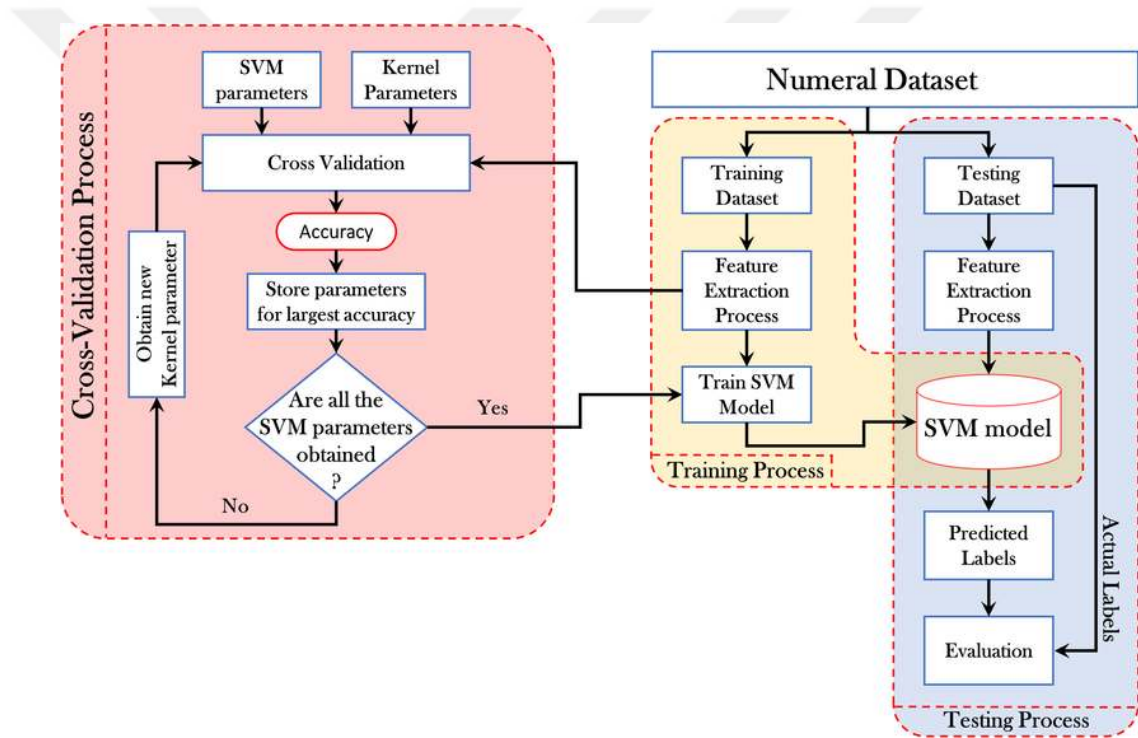


Figure 3.2 A schematic diagram for support vector machine (SVM) training and testing process.

Utilizing the SVM method is justified by its ability to generate a hyperplane and maximize the margin between the data and separation classes of the hyperplane. Through out-of-sample error control, SVM can also reduce structural risk. Furthermore, SVM is better suited for recognition because it can withstand a noisy environment better. The positive and negative images can be distinguished by SVM by creating a hyperplane. The SVM method's classification procedure, including the

training and testing phases, is schematically depicted in Figure 3.2. Both polynomial and radial basis functions (RBF) are used by the SVM kernels. These kernels exhibit nonlinear separation between classes, making them an effective classification mechanism. High prediction accuracy is ensured by the cross-validation process. The optimal kernel parameters can thus be achieved (Abdulhussain et al., 2021).

In conclusion, SVM offers a powerful algorithm for linear and nonlinear classifications. The use of kernel functions allows data to be separated by more complex boundaries, while maintaining the widest possible margin increases the model's accuracy. These features make SVM an effective tool for solving many different machine learning problems, such as classification and regression.

3.3 The Methodology of GEP

Gene Expression Programming, as a type of evolutionary algorithm, offers a method that allows a solution to evolve through the evolutionary process using evolutionary principles. GEP works similarly to genetic algorithms and genetic programming, but unlike them, it has a more flexible and advanced structure. The operation of GEP consists of the following basic steps:

Structure of Chromosomes: GEP represents genetic codes as a type of chromosome structure. These chromosomes contain genetic programs that express solution proposals. Each chromosome consists of a combination of genetic expressions and functions, and these genetic expressions create a solution for a specific problem. The structural organization of chromosomes mimics the functions and operators in a programming language.

Evolutionary Process and Selection: GEP is based on the fundamental principles of evolutionary algorithms. Solution proposals evolve through evolutionary operators such as selection, crossover, and mutation. The crossover process allows two chromosomes to combine and produce a new solution, while the mutation process simulates random changes in the chromosomes. These evolutionary processes enable the production of better solutions over time.

Objective Function and Optimization: The main goal of GEP is to find the most suitable solution for a specific problem. This process is achieved by optimizing the objective function. The objective function is used to determine the best solution related to the problem, and the evolutionary process aims to achieve the minimum or maximum value of this function. GEP aims to reduce the error rate in the solution by using genetic structures and functions.

Evaluation of Results: The advanced structure of GEP allows it to achieve high accuracy on solutions. The accuracy of the solution is generally measured through evaluations conducted on test data. These accuracy rates indicate the model's success and lead to the pursuit of better solutions.

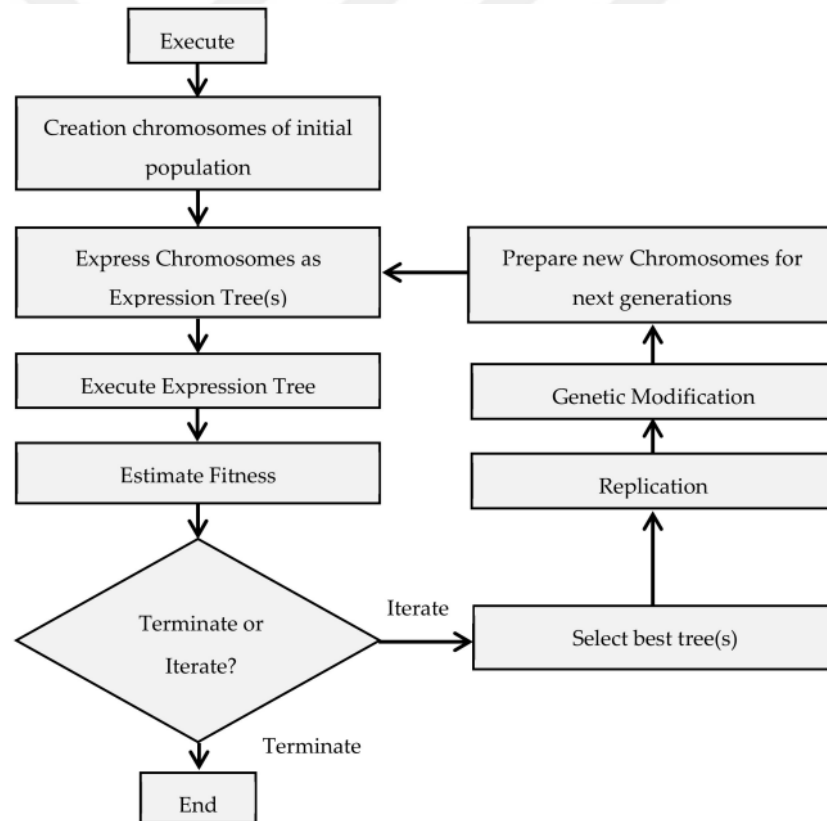


Figure 3.3 An illustration of Gene Expression Programming (GEP) algorithm

To sum up, GEP is an effective tool for solution generation and function discovery because it combines genetic algorithms and programming. Its ability to resolve intricate and nonlinear issues, like hydrological modeling, makes it useful for a variety of applications.

3.4 The Methodology of Ridge Regression

Ridge regression is a type of linear regression that uses a regularization technique to prevent the parameters of the model from becoming excessively large. This method is particularly preferred in situations where there are multiple linear relationships and the problem of multicollinearity (high correlation between variables) is encountered. Ridge regression adds a penalty term to the model's loss function. This penalty term is proportional to the sum of the squares of the model's weights, and thus, by reducing the model's parameters, overfitting is prevented. Mathematically, ridge regression can be expressed as follows:

$$\hat{\beta} = \arg \min_{\beta} \left[\sum_{i=1}^n (y_i - X_i \beta)^2 + \lambda \sum_{j=1}^p \beta_j^2 \right]$$

Here, λ is the regularization parameter, and the larger it is, the more the model's parameters are shrunk. An important feature of ridge regression is that it shrinks all regression coefficients but does not reduce them to zero, which distinguishes it from lasso regression. While lasso performs variable selection by setting some coefficients to zero, ridge regression does not make such a selection. The applications of ridge regression are found in a wide range of fields such as financial models, biomedical analyses, and machine learning (Hoerl and Kennard, 1970). The strength of ridge regression lies in its ability to enhance the model's predictive capability and make it more generalizable (Tikhonov, 1943).

3.5 Shapley Additive Explanations (SHAP) and Its Use with Machine Learning

SHAP (Shapley Additive Explanations) analysis is a powerful method aimed at enhancing the interpretability of machine learning (ML) model predictions. It takes its name from the Shapley values used in cooperative game theory. Shapley values are a theoretical framework that fairly calculates the contribution of a player (or feature). SHAP analysis adapts this theory to ML models, determining the contribution of each feature to the model's prediction. In this way, the inner workings of complex models, often referred to as "black boxes," become more transparent. In machine learning, SHAP clearly reveals how each feature affects the model's

prediction. It is particularly used in complex models such as decision trees, random forests, and deep learning (Lundberg and Lee, 2017). SHAP explains the effect of each feature on the prediction, whether positive or negative. The contributions of these features allow for a better understanding of the model's overall performance and enhance the model's explainability.

SHAP values show the contribution of each feature to the model's output. If a feature's SHAP value is positive, it has contributed to an increase in the prediction; if negative, it has contributed to a decrease in the prediction. SHAP can analyze the impact of each feature both globally (overall) and locally (for a single prediction). Global analysis reveals the overall behavior of the model, while local analysis shows how a single prediction is made. The visual presentation of SHAP values is particularly useful in understanding the impact magnitudes of features (Ribeiro et al., 2016).

3.6 Case Study for Imputation Missing Data

In this study, D20A036 flow gauge station's monthly discharge data, which includes observed data between 1979 and 2014 years, was used. Between the years 1979-2014, the data for the year 1989 and between the years 1997-2010 are missing. In this study, the data for the mentioned missing years will be estimated. The station is located in the village of Sisne where is 28 km away from Andırın district, Kahramanraş Türkiye. The local station has a latitude (N) of $37^{\circ}43'$ N, longitude (E) of $36^{\circ}29'$.

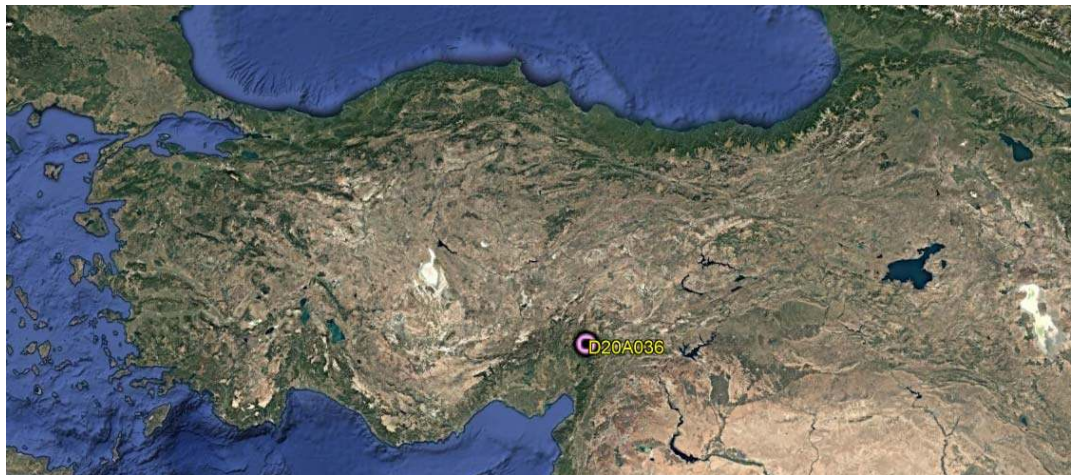


Figure 3.4 Location of D20A036 flow gauge station

Table 3.2 Station discharges (m³/s) measure results with missing data

Station Name		D20A036 Kösulu River, SİSNE										
Location		In the vilage of Sisne where is 28 km away from Andırın district. (36°29' E-37°43' N)										
Rain Area		150.8 km ²					Altitude			1250 m.		
Year	Oct	Nov	Dec	Jan	Feb	Marc	Apr	Ma	Jun.	July	Aug	Sept
1979	0.74	1.75	18.8	31.5	26.3	25	28.	25	17.8	14.3	12.1	11.1
1980	3.3	1.01	2.21	1.29	1.97	23.3	34.	23	3.77	1.3	0.67	0.30
1981	0.40	0.37	1.38	2.15	2.94	22.7	15.	11.	4.42	1.39	0.62	0.34
1982	0.64	0.84	10.8	10.1	2.39	4.29	18.	10	2.72	1.06	0.45	0.36
1983	0.44	0.36	0.16	0.15	0.20	6.94	16.	8.9	2.76	0.92	0.42	0.24
1984	0.41	4.41	3.23	1.94	7.1	8.59	12	8.0	1.94	0.76	0.33	0.13
1985	0.16	0.33	0.16	0.57	3.57	6.18	11.	4.5	1.17	0.14	0.01	0
1986	1.09	1.15	1	6.1	2.72	5.3	7.0	3.8	2.17	0.71	0.35	0.12
1987	0.23	0.71	1.47	6.31	8.93	6.05	31.	30.	8.39	2.03	0.91	0.25
1988	0.37	1.15	12.7	2.28	4.7	16.5	26.	17.	4.6	1.58	0.48	0.27
1989	--	--	--	--	--	--	--	--	--	--	--	--
1990	1.99	10.7	5.7	2.15	2.05	4.37	12.	7.2	2.55	1	0.49	1.8
1991	1.16	1.72	2.01	0.68	0.88	7.99	13.	5.6	2.17	0.67	0.28	0.11
1992	0.44	1.76	2.09	1.51	2.14	5.79	26.	17.	4.67	1.22	0.36	0.11
1993	0.13	0.68	1.12	0.72	1.58	6.24	21.	14.	4.69	0.92	0.30	0.04
1994	0.08	0.18	0.29	2.07	1.76	6.62	7.4	4.6	1.11	0.26	0.04	0
1995	0.13	3.92	2.01	6.43	4.21	7.6	18	10.	3.15	1.12	0.26	0.02
1996	0.37	7.99	2.94	7	6.13	14.8	25.	17.	3.97	0.87	0.33	0.91
1997	--	--	--	--	--	--	--	--	--	--	--	--
1998	--	--	--	--	--	--	--	--	--	--	--	--
1999	--	--	--	--	--	--	--	--	--	--	--	--
2000	--	--	--	--	--	--	--	--	--	--	--	--
2001	--	--	--	--	--	--	--	--	--	--	--	--
2002	--	--	--	--	--	--	--	--	--	--	--	--
2003	--	--	--	--	--	--	--	--	--	--	--	--
2004	--	--	--	--	--	--	--	--	--	--	--	--
2005	--	--	--	--	--	--	--	--	--	--	--	--
2006	--	--	--	--	--	--	--	--	--	--	--	--
2007	--	--	--	--	--	--	--	--	--	--	--	--
2008	--	--	--	--	--	--	--	--	--	--	--	--
2009	--	--	--	--	--	--	--	--	--	--	--	--
2010	--	--	--	--	--	--	--	--	--	--	--	--
2011	0.47	0.30	1.65	1.26	3	10.6	24.	10.	2.59	0.91	0.29	0.22
2012	0.32	0.36	0.45	1.68	1.06	5.01	43.	18.	3.04	0.70	0.14	0.1
2013	0.29	1.2	3.82	3.93	5.52	9.63	10.	5.1	1.98	0.30	0.08	0.05
2014	0.20	0.16	0.16	0.30	0.40	4.52	2.2	2.0	0.57	0.07	0.01	0.01

3.7 Case Study for Comparison Methods and Projections

For the training period, data from 1988 to 2006 were used, and for the testing period, data from 2007 to 2014. Throughout the training phase, historical climate data is used to calibrate and optimize the model's parameters, assisting it in finding trends and connections in the data. Along with more conventional techniques like GEP, SVM, Ridge Regression, and the new boosting model CatBoost, the data, which was separated into training and testing sets, was examined using historical data from the NorESM2-MM model for the years 1979–2014. Following the selection of the strategy with the best test model performance, projections were made under a number of future climate change scenarios. The decision-making process flow diagram for the model's performance is displayed in the Figure 3.5.

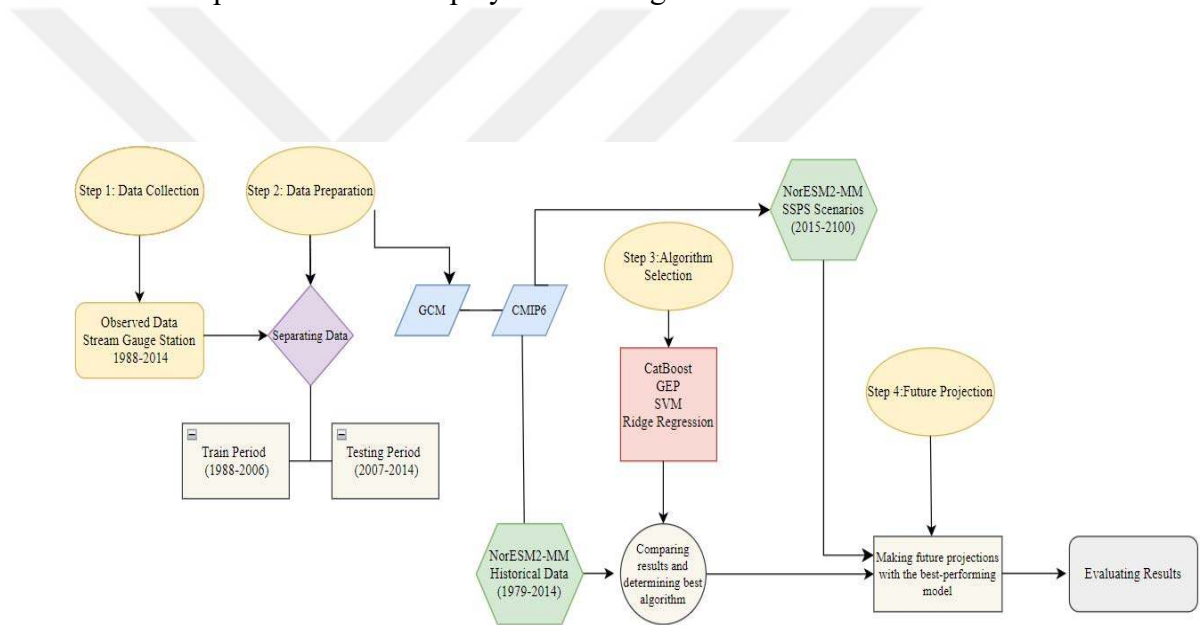


Figure 3.5 Flow diagram of the decision process for model performance and future projections

Table 3.3 Station informations and historical discharges (m³/s)

Station Information												
Station Name			D20A053 Fırnız River Kabaktepe									
Station Location			It is located at the 52nd kilometer of the Kahramanmaraş-Göksun Road, 20 kilometers downstream of the new Fırnız Bridge. 36°41' D-37°45' K									
Rainfall Area			147.1 km ²				Altitude			648 m.		
Year	Oct	Nov	Dec	Jan	Feb	Mar.	Apr.	Ma	Jun	Jul	Aug	Sept
1988	1.91	2.44	10.1	5.88	8.46	19.6	30.8	21.2	7.08	3.36	2.77	2.27
1989	2.38	3.55	7.05	4.16	3.77	7.75	5.5	3.51	2.11	1.9	1.64	1.59
1990	2.76	8.64	19.7	6.73	7.72	10.5	11.8	8.39	4.15	3.19	2.32	1.99
1991	1.77	2.44	4.77	4.45	5.79	11.2	16.2	9.65	4.06	2.59	2.12	1.85
1992	1.72	1.72	4.79	4.35	4.12	8.6	23.7	16.4	6.72	3.17	2.46	1.93
1993	1.36	1.91	3.59	3.03	3.65	9.29	22.3	15.8	6.57	3.12	2.31	1.82
1994	2.12	1.87	1.82	3.81	4.63	8.1	8.92	5.04	2.51	2.14	1.86	1.63
1995	1.48	4.61	3.86	9.5	7.42	11.3	17.5	10.5	5.5	3.09	2.66	2.12
1996	1.54	6.02	3.86	11.7	10.5	9.79	16.1	16.8	6.33	3.27	2.47	2.51
1997	3.19	2.68	8.41	8.37	4.97	4.47	15.3	11.4	4.2	2.4	1.95	1.7
1998	1.55	2.28	6.09	3.84	5.77	9	17.4	12.6	5.15	3.08	2.2	1.69
1999	1.52	1.74	9.04	5.69	10.1	9.87	13.3	6.36	3.26	2.55	1.83	1.7
2000	1.23	1.17	1.23	2.36	3.56	8.15	22.6	10.3	3.88	2.36	1.99	1.73
2001	1.63	1.45	1.4	1.39	1.78	7.89	6.65	5.9	2.96	2.09	1.78	1.62
2002	1.48	1.34	7.68	5.41	4.42	15.5	25.6	14.7	5.21	3.54	2.78	2.22
2003	2.16	1.74	1.54	3.76	6.13	8.77	23.7	10.4	6.41	3.26	2.71	2.25
2004	1.86	1.62	3.2	6.11	6.82	13.9	12.5	7.47	2.65	2.34	2.01	1.82
2005	1.59	2.34	1.97	2	3.78	19.3	15.9	7.49	3.33	2.23	2.01	1.63
2006	1.73	1.85	3.68	2.5	6.01	11.6	12.1	4.71	2.41	2.02	1.77	1.67
2007	1.83	4.08	1.65	1.34	2.32	8.72	6.79	5.74	3.32	2.07	1.67	1.26
2008	1.42	2.07	5.82	1.99	2.32	12.8	10.6	3.45	2.22	1.92	1.57	1.42
2009	1.81	3.07	2.17	3.36	14.5	20.4	32.8	19.9	5.03	2.44	1.83	1.67
2010	1.77	2.44	4.77	4.45	5.79	11.2	16.2	9.65	4.06	2.59	2.12	1.85
2011	1.5	1.11	4.29	2.55	5.11	13.8	18.8	10.9	4.4	2.82	2.38	1.83
2012	1.76	1.74	1.77	4.17	3.36	10	30.7	13.9	4.16	2.41	1.89	1.77
2013	1.96	2.25	10.9	6.74	11.8	12.8	12.5	6.21	3.1	2.03	1.84	1.68
2014	1.58	1.36	1.32	1.33	1.39	5.73	2.34	1.75	1.7	1.46	1.86	1.86

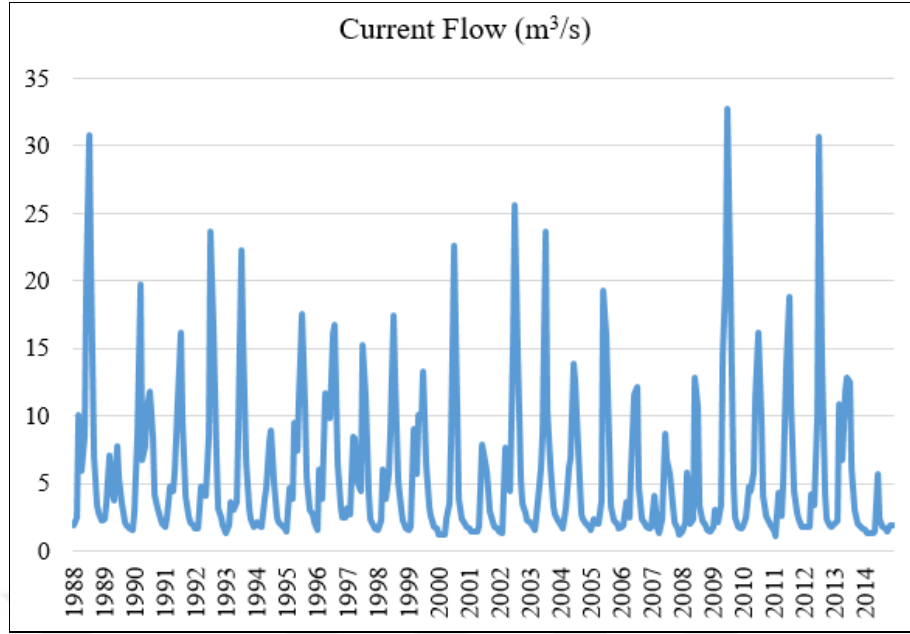


Figure 3.6 Mean monthly discharge amounts between 1988 and 2014

Data from the years 1988 to 2006 was employed for the training period, while data from the years 2007 to 2014 was used for the testing period. During the training period, the model's parameters are calibrated and optimized using past climate data, helping it to discover patterns and relationships within the data. The data divided into training and testing sets was analyzed using historical data from the NorESM2-MM model for the years 1979-2014 along with traditional methods such as GEP, SVM, Ridge Regression and the new boosting model CatBoost.

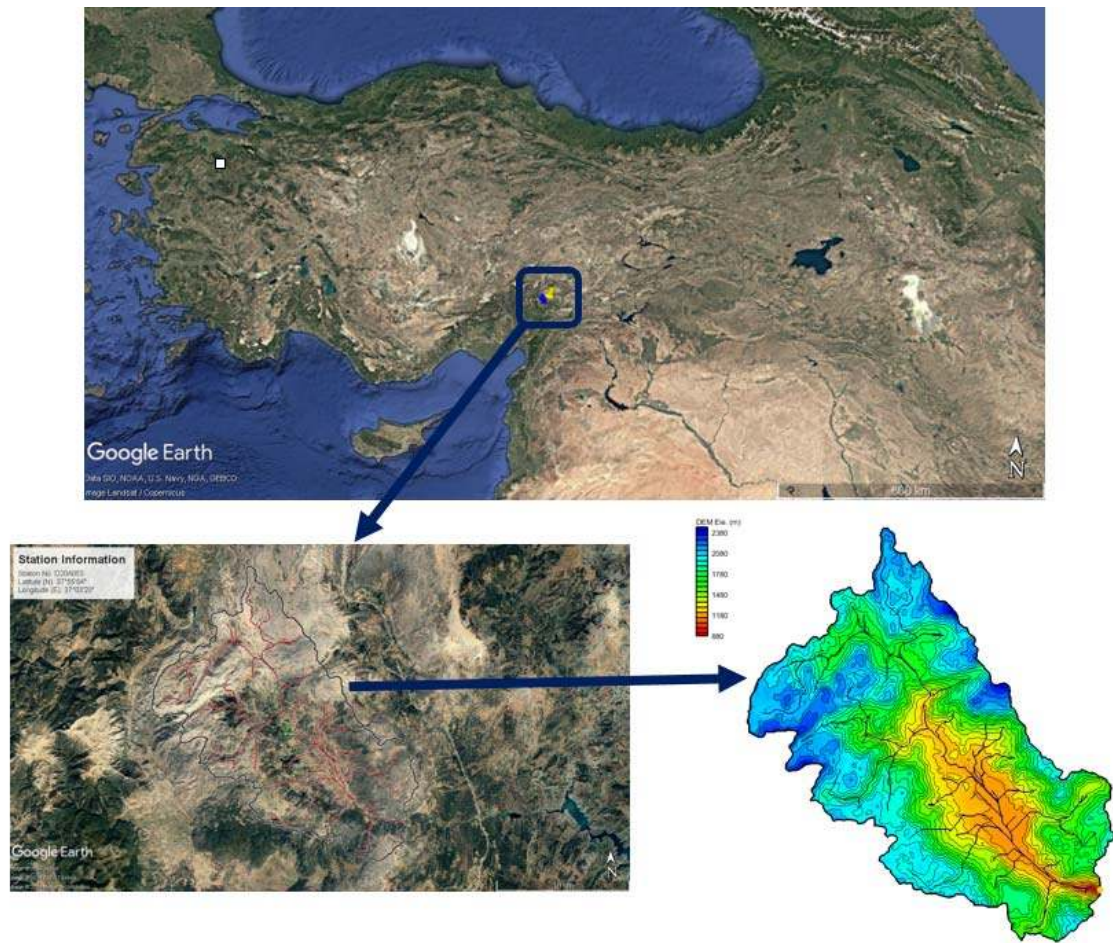


Figure 3.7 Location of Gauge Station and Digital Elevation Model (DEM) contours of the watershed.

3.8 Handling NorESM2 Data

As part of the sixth phase of the Coupled Model Intercomparison Project (CMIP6), the Norwegian Earth System Model version 2 (NorESM2) is being prepared. By offering a more thorough interaction between the atmosphere, ocean, land, and ice components at various resolution levels, the NorESM2 model has been chosen because it can replicate the intricate dynamics of the climate system and highlight local climate features. There are two versions of the model: NorESM2-LM, which is low-resolution, and NorESM2-MM, which is medium-resolution. Predictor factors were incorporated into the ensemble using NorESM2-MM, which has a higher resolution. The datasets are based on a 64 by 128 latitude-longitude global Gaussian grid with a spectral truncation of T42 and have atmospheric horizontal resolution. The native grid of the CanESM5 has a roughly uniform latitudinal resolution of

2.8125° and a consistent longitudinal resolution of 2.8125°. To conform to the CanESM5 grid, NorESM2 data were interpolated.

Files called BOX_iiiX_jjY, where iii is the longitudinal index and jj is the latitudinal index, contain the predictions for every grid cell. After selecting the appropriate grid cell, the data set is downloaded as a zip file from the Canadian Climate Data and Scenarios website (<https://climate-scenarios.canada.ca/?page=pred-cmip6#cmip6-predictors>). Either manually entering the stream's center or choosing the cell from the page's rectangle map will accomplish this. The zip file 'BOX_014X_46Y' was obtained for the designated grid cell by clicking the retrieve data button. 26 variables are included in the dataset for the following four scenarios: SSP1-2.6, SSP2-4.5, SSP3-7.0, and SSP5-8.5. The data set of daily predictors was transformed into monthly averages. The variables of NorESM2 are described in Table 3.4.

Table 3.4 List of the NorESM2 variables (predictors) IDs and corresponding variable names

No.	Variable ID	Predictor variable
1	mslp	Mean sea level pressure
2	p1_f	1000 hPa Wind speed
3	p1_u	1000 hPa Zonal wind component
4	p1_v	1000 hPa Meridional wind component
5	p1_z	1000 hPa Relative vorticity of true wind
6	p1th	1000 hPa Wind direction
7	p1zh	1000 hPa Divergence of true wind
8	p5_f	500 hPa Wind speed
9	p5_u	500 hPa Zonal wind component
10	p5_v	500 hPa Meridional wind component
11	p5_z	500 hPa Relative vorticity of true wind
12	p5th	500 hPa Wind direction
13	p5zh	500 hPa Divergence of true wind
14	p8_f	850 hPa Wind Speed
15	p8_u	850 hPa Zonal wind component
16	p8_v	850 hPa Meridional wind component

17	p8_z	850 hPa Relative vorticity of true wind
18	p8th	850 hPa Wind direction
19	p8zh	850 hPa Divergence of true wind
20	p500	500 hPa Geopotential
21	p850	850 hPa Geopotential
22	prcp	Total precipitation
23	s500	500 hPa Specific humidity
24	s850	850 hPa Specific humidity
25	shum	1000 hPa Specific humidity
26	temp	Air temperature at 2 m

3.9 Downscaling using GEP and SVM techniques

For GEP, this study highlights that choosing a set of functions to be used is the first step. Determining the chromosome structure, which includes defining the quantity and size of genes, is the second step. Choosing the linking function is the third step. Assessing fitness with a particular metric is the last phase. The foundation of GEP consists of these four fundamental steps, which must be carefully followed in order to produce useful results in problem-solving applications. The GEP model's main purpose is to generate a mathematical equation from training data (Fuladipanah et al., 2023).

To find the best answers, this method iteratively improves the population, simulating natural evolution (Ferreira, 2001).

The evolution of solutions in GEP is influenced by hyperparameters such as population size and mutation rate, which have an impact on model accuracy and convergence. The population size=20000, Generations=40, Subtree mutation=0.1 and Persimony coefficient=0.01 of the GEP are the hyperparameters used in this study.

Support Vector Machines's (SVMs) first step is to prepare the data by splitting it into training and test sets. Missing values and normalization are also completed in this step. This is particularly important in SVM because the algorithm relies on

calculating the distances between data points in high-dimensional space. Then, it uses kernels to transform the data into a higher-dimensional space. The most popular kernels are linear, polynomial, and Radial Basis Function (RBF) kernels; the choice of kernel depends on the problem and the type of data, with RBF working best for non-linear classification problems. Then, the model is trained by determining the best hyperplane in the dataset that maximizes the margin between the various classes. After the model is trained, its performance is evaluated using test data. SVM's capacity to generalize and prevent overfitting is governed by parameters like the kernel type, regularization, and margin. The kernel, gamma, epsilon=1, and regularization parameter (C)=40 of the SVM are the hyperparameters used in this study.

The efficacy of SVM is rooted in its capacity to identify a decision boundary that exhibits good generalization, even in intricate and high-dimensional spaces (Cortes & Vapnik, 1995).

3.10 Assessment Metrics

Metrics notably Root Mean Squared Error (RMSE), Mean Absolute Error (MAE), Kling-Gupta Efficiency (KGE), RSR (Root Mean Square Error to Standard Deviation Ratio) and Nash-Sutcliffe Efficiency (NSE) offer unique insights into prediction model accuracy and performance.

RMSE estimates the average size of prediction errors by using the square root of the average squared differences between predicted and observed values, yielding an error measure in the same units as the data (Hyndman & Athanasopoulos 2018). The square of RMSE, emphasizes the average of squared differences without unit normalization, providing greater weight to larger errors (Montgomery et al., 2012). Mean Absolute Error (MAE) is a statistic that assesses the accuracy of a model's predictions by calculating the average absolute difference between projected and observed values. It provides a simple interpretation of a model's average error, making it a popular choice for regression analysis and forecasting.

To evaluate model performance the root mean square error (RMSE) and mean absolute error (MAE) are calculated. A ideal model would have an RMSE and MAE of 0. Both errors are specified in the same units as the predicted quantities.

The Nash-Sutcliffe model Efficiency (NSE), a commonly used statistic to assess the performance of hydrologic models, is calculated. Nash and Sutcliffe (1970) introduced NSE, which evaluates model performance relative to a mean value model, with values ranging from $-\infty$ to 1, where 1 indicates a perfect model fit and values below zero indicate that the model performs worse than using the mean of observed values as a predictor (Moriassi et al., 2007).

RMSE and MAE quantify error size, while NSE compares model efficiency, which is significant in domains like hydrology where performance versus a simple mean is critical (Krause et al., 2005). In this study, RMSE, MAE, and NSE will be calculated and taken into account to measure the performance of the methods used in calculating the estimated mean discharge value produced for the future using historical data under climate change scenarios.

The Kling-Gupta Efficiency (KGE) metric is used for assessing the performance of hydrological models, especially when comparing simulated streamflow to observed data. By taking into account the relative bias and variability as well as the correlation between the simulated and observed time series, KGE aggregates several aspects of model performance into a single value. Three elements make up the KGE formula: the variability ratio (γ), the bias ratio (β), and the correlation coefficient (r). In the KGE metric, which has a range of $-\infty$ to 1, a value of 1 denotes perfect agreement between the observed and simulated data, whereas values nearer 0 or negative denote subpar model performance. KGE evaluates the performance of machine learning algorithms (Mishra et al., 2024).

In hydrological prediction, the RSR (Root Mean Square Error to Standard Deviation Ratio) metric is frequently used to evaluate model performance, especially when forecasting rainfall, river discharge, and other hydrological variables. It is a dimensionless ratio that contrasts the observed data's inherent variability with the model's error magnitude. It is calculated mathematically by dividing the standard deviation of the observed data by the root mean square error (RMSE) of the model predictions. Better model performance is indicated by a lower RSR value; in general, a good model fit is indicated by an RSR of less than 0.5. Because it normalizes the error in relation to the variability in the observed data, the RSR metric is useful for

comparing various datasets with varying scales and magnitudes (Moriassi et al., 2007).

3.10.1 Root Mean Squared Error (RMSE)

Root Mean Squared Error (RMSE) is a performance evaluation metric that measures the average squared error between the values predicted by the model and the actual observed values. RMSE is frequently used to understand the magnitude of errors in the model's predictions. When measuring the magnitude of errors, it squares the errors, thereby ensuring that larger errors are penalized more. This helps to evaluate how precise the model is.

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_{\text{obs}}(i) - y_{\text{sim}}(i))^2}$$

$y_{\text{obs}}(i)$: Observed value (i. observation).

$y_{\text{sim}}(i)$: The value predicted by the model (i observation).

n : Number of observations.

Method of Calculation:

- For each observation, the square of the difference between the prediction and the actual value is taken.
- The squares of all observation differences are summed.
- This total is divided by the number of observations (n).
- Finally, the RMSE value is obtained by taking the square root of this average value.

Interpretation:

- $RMSE = 0$: The model's predictions fit perfectly.
- $RMSE > 0$: The model's predictions are inconsistent with the observations.
- The smaller the RMSE value, the better the model's predictions.

3.10.2 Mean Absolute Error (MAE):

Mean Absolute Error (MAE) is a statistic that measures the average of the absolute differences between the model's predictions and the actual observed values. MAE

measures the magnitude of errors more directly by taking the absolute value of each error. Compared to RMSE, MAE penalizes large errors less because it uses absolute differences instead of squared differences.

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_{\text{obs}}(i) - y_{\text{sim}}(i)|$$

$y_{\text{obs}}(i)$: Observed value (i. observation).

$y_{\text{sim}}(i)$: The value predicted by the model (i observation).

n : Number of observations.

Method of Calculation

- For each observation, the absolute value of the difference between the prediction and the observation is taken.
- All absolute differences are summed up.
- This total is divided by the number of observations (n) to calculate the MAE value.

Interpretation:

- $MAE = 0$: The model's predictions fit perfectly.
 - $MAE > 0$: The model's predictions are inconsistent with the observations.
- The smaller the MAE value, the better the model's predictions.

Differences Between RMSE and MAE:

- RMSE penalizes larger errors more (because the error is squared), while MAE treats all errors equally.
- RMSE is a more sensitive measure of error magnitude and increases the impact of large errors. MAE, on the other hand, takes the average of the errors and is a more stable measure.
- RMSE is generally preferred in modeling applications that require greater precision, while MAE provides a simpler and more understandable evaluation (Willmott and Matsuura, 2005).

3.10.3 Root Mean Square Error to Standard Deviation Ratio (RSR)

Root Mean Square Error to Standard Deviation Ratio (RSR) is a performance evaluation metric that measures the ratio of the model's prediction errors to the

standard deviation of the observed data. This ratio takes into account how variable the observed data is when evaluating the accuracy of the model's predictions.

RSR is particularly used to understand how good the model is because it allows for the comparison of the model's prediction errors with the variance of the observed data.

RSR is generally used in hydrological modeling and time series forecasting, and it compares the magnitude of the model's error with the standard deviation of the observed data. This is useful for evaluating the model's accuracy against a reference value.

$$RSR = \frac{RMSE}{\sigma_{obs}}$$

- RMSE: Root Mean Squared Error
- σ_{obs} : The standard deviation of the observed data.

Method of Calculation

- First, the RMSE value is found by calculating the error between the model's predictions and the observations.
- The standard deviation of the observed data is calculating.
- The RMSE value is divided by the standard deviation of the observed data to obtain the RSR value.

Interpretation:

- $RSR = 0$: The model makes a perfect prediction because the error is zero.
- $RSR > 1$: The model's prediction errors are large compared to the observations, and the model's accuracy is low.
- $RSR < 1$: The model's predictions have been made with quite high accuracy and low errors compared to the observed data.

The RSR value provides a simple and effective way to determine the success of the model. However, cases where the RSR is less than 1 are generally considered a good model.

3.10.4 Nash-Sutcliffe Model Efficiency (NSE)

Nash-Sutcliffe Model Efficiency (NSE) is a statistical measure that evaluates the accuracy of a model in hydrological modeling and simulations. This metric shows how well the model aligns with the observations. The NSE value evaluates how successful the model is by comparing the model's predictions with the observations (Nash and Sutcliffe, 1970). The mathematical definition of NSE is as follows:

$$NSE = 1 - \frac{\sum_{t=1}^n (Q_{\text{obs}}(t) - Q_{\text{sim}}(t))^2}{\sum_{t=1}^n (Q_{\text{obs}}(t) - \overline{Q_{\text{obs}}})^2}$$

Here:

$Q_{\text{obs}}(t)$: The observed flow value at time t in the time series.

$Q_{\text{sim}}(t)$: The value of the flow rate (simulated flow) predicted by the model at time t in the time series.

$\overline{Q_{\text{obs}}}$: The average of the observed flow.

n : Number of observation points (time steps).

Method of Calculation:

- Observed Values and Predictions: First, the values predicted by the model $Q_{\text{sim}}(t)$ and the observations $Q_{\text{obs}}(t)$ should be taken.
- Squares of the Differences from the Mean of Observations: The mean of the observed values is calculated, and the square of the difference of each observation from this mean is taken.
- Sum of the Squares of the Differences Between Prediction and Observation: For each time period, the square of the difference between the observed value and the predicted value is taken.
- NSE Calculation: The sum of the squares of these differences is subtracted from the sum of the squares of the differences between the observations and the mean, and the result is subtracted from 1 to obtain the NSE value.

NSE Interpretation:

- $NSE = 1$: The model shows perfect agreement (the model's predictions are fully consistent with the observations).

- $0 < \text{NSE} < 1$: The model's accuracy is acceptable, but the model's predictions deviate somewhat from the observations.
- $\text{NSE} = 0$: The accuracy of the model is as good as a simple average model.
- $\text{NSE} < 0$: The model yields worse results than a random guess.

3.10.5 Kling-Gupta Efficiency (KGE)

Kling-Gupta Efficiency (KGE) is another performance evaluation metric used in hydrological modeling and simulations. KGE evaluates the accuracy of the model through three main components: amplitude, timing, and correlation. This measures not only how well the model aligns with observations but also how accurately the model predicts the dynamics and amplitudes over time. KGE is generally used to evaluate the success of a model, like the Nash-Sutcliffe Model Efficiency (NSE), but KGE is more comprehensive because it takes into account not only the magnitude of the errors but also the structural aspects of the errors (Gupta et al., 2009).

$$KGE = 1 - \sqrt{(r - 1)^2 + \left(\frac{\sigma_{\text{sim}}}{\sigma_{\text{obs}}} - 1\right)^2 + \left(\frac{\mu_{\text{sim}}}{\mu_{\text{obs}}} - 1\right)^2}$$

r : The correlation coefficient between the model and the observations (r , Pearson correlation).

σ_{sim} : The standard deviation of the model's predictions.

σ_{obs} : The standard deviation of the observed data.

μ_{sim} : The average of the model's predictions.

μ_{obs} : The average of the observed data.

Method of Calculation:

- Correlation (r): First, the Pearson correlation coefficient between the model's predictions and the observations is calculated. This shows how similar the time series are.
- Standard Deviations (σ): The standard deviations of both the model's predictions and the observed data are calculated. This is used to evaluate amplitude differences.
- Average Values (μ): The averages of the model's predictions and the observed data are calculated. This is used to evaluate the model's amplitude errors.

- **KGE Calculation:** According to the KGE formula, the square of the deviations of all three components (correlation, amplitude, and time) is taken, and the square root of these is subtracted from 1.

KGE Interpretation:

- $KGE = 1$: Perfect model. The model perfectly fits the observations.
- $KGE = 0$: Acceptable model. The model fits the observations with some errors.
- $KGE < 0$: The model makes completely random predictions and fits poorly with the observations.

In this study, historical data obtained from the NorESM2-MM model, which is part of CMIP6, were analyzed using traditional methods such as Ridge Regression, GEP, and SVM, as well as an innovative method called CatBoost, and the accuracy of each model was tested. When determining accuracy, metrics such as NSE, RMSE, MAE, KGE, and RSR have been taken into account. Later, the method that yielded the best results according to the test period was identified, and future projections were developed based on this method with 26 inputs.

3.11 Mann-Kendall Test

The Mann-Kendall test is a statistical test used to determine trends in time series data. This test is used to evaluate whether the data shows a trend in a specific direction (increasing or decreasing) and is widely applied, especially in environmental sciences, hydrology, and climate change analyses. The Mann-Kendall test does not require the data to be ordered or the assumptions of normal distribution in the data. Therefore, it is considered a non-parametric test. Purpose of the test is checks whether the data in a time series shows a linear trend (Kendall,1975). The trend can be increasing (positive) or decreasing (negative). The Mann-Kendall test compares the differences between each data point in the dataset and the other data points. The basic formula of this test is as follows:

Dataset: When there is a time series data in the form of $x_1, x_2, \dots, x_n$, the differences between each data point x_i and x_j ($i < j$) are calculated.

Count-Based Calculation (C):The Mann-Kendall test uses the number of signs obtained through the ranking process. For each pair of x_i and x_j , if $x_j > x_i$, a positive sign (+1) is given to increase the SSS value, and if $x_j < x_i$, a negative sign (-1) is given to decrease the SSS value. For observations that have the same value $x_i = x_j$ no sign is given.

$$S = \sum_{i=1}^{n-1} \sum_{j=i+1}^n \text{sgn}(x_j - x_i)$$

Test Statistic (Z):To ensure that the S value approaches a normal distribution, the test's z-score is calculated as follows:

The mean $E(S)$ and variance $Var(S)$ are calculated.

$$E(S) = 0$$

$$Var(S) = \frac{n(n-1)(2n+5)}{18}$$

If the length of the series is small, corrections can be made.

Z-score is calculated:

$$Z = \frac{S - 1}{\sqrt{Var(S)}}$$

$$Z = \frac{S + 1}{\sqrt{Var(S)}}$$

Method of Calculation:

- Data sorting: First, the data in the time series is sorted, and the differences between each pair of data points are calculated.
- Calculation of signs: The signs of the differences between x_i and x_j are determined (+1, -1, 0).
- Calculation of total signs: The sum of positive and negative signs, known as SSS, is calculated.
- Calculation of the Z-score: The Z-score is calculated using the value of S.

- Interpretation of the results: The Z-score is evaluated against a specific criterion (generally $|Z| \geq 1.96$ when at a 95% confidence level) to determine whether there is a trend.

Interpretation:

- Positive Z: If the Z score is positive, it indicates an increasing trend in the data.
- Negative Z: If the Z score is negative, it indicates a decreasing trend in the data.
- $Z = 0$: If the Z score is zero, it is concluded that there is no trend.
- $|Z| \geq 1.96$: There is a statistically significant trend (95% confidence level).

3.12 Prediction Energy Production

The formula for calculating hydropower production is:

$$P = \eta \times \rho \times g \times Q \times H \quad (3)$$

Where:

- P = Power output (in watts, W)
- η = Overall Efficiency of the turbine and generator (dimensionless, typically between 0.7 and 0.9)
- ρ = Density of water (approximately 1000 kg/m³ at 4°C)
- g = Acceleration due to gravity (9.81 m/s²)
- Q = Flow rate of water (in cubic meters per second, m³/s)
- H = Height of the water head (in meters, m)

This formula gives the potential power that can be generated from a hydropower plant, assuming all other factors like efficiency are ideal. The following values have been used to observe the impact of future flow values on the amount of energy to be produced; no dam-related data will be used in the energy calculation for this study.

- $\eta = 0.8$
- $\rho = 1000 \text{ kg/m}^3$
- $g = 9.81 \text{ m/s}^2$

While calculating the impact of the obtained flows on energy production, many complex data such as evaporation, HDD and CDD values, spillway quantity, dam volume, etc., were not taken into account.

Although it is possible to perform more complex calculations for a more realistic and reliable result, in this study, only flow-related changes were predicted at a simpler level.

The dam heights to be used in this calculation were determined by mass curve analysis. A mass curve analysis using the anticipated flows and cumulative values was used to determine the expected volume for the each scenarios between years 2015–2100. Given that the flow observation station in question—for which future discharge is anticipated due to climate change—is situated in the Ceyhan basin and that the Menzelet Dam is the closest dam impacted by the measured values at the station, the height produced by the flows was computed by dividing the volumes derived from the mass curve by the Menzelet Dam's present lake area 42 km².

Table 3.5 Calculated dam heights results under each scenario according to mass curve analysis

	Volume (m ³)	Area (m ²)	Height (m)
Observed	113,1414	42000000	2,69
SSP1-2.6	65,0839	42000000	1,55
SSP2-4.5	89,841	42000000	2,14
SSP3-7.0	169,483	42000000	4,04
SSP5-8.5	211,301	42000000	5,03

The station and Menzelet Dam's location given in Figure 3.8. The anticipated average flow and energy production amounts are displayed in Table 4.10. Furthermore, Figure 4.38 compares the anticipated energy amounts to be generated.

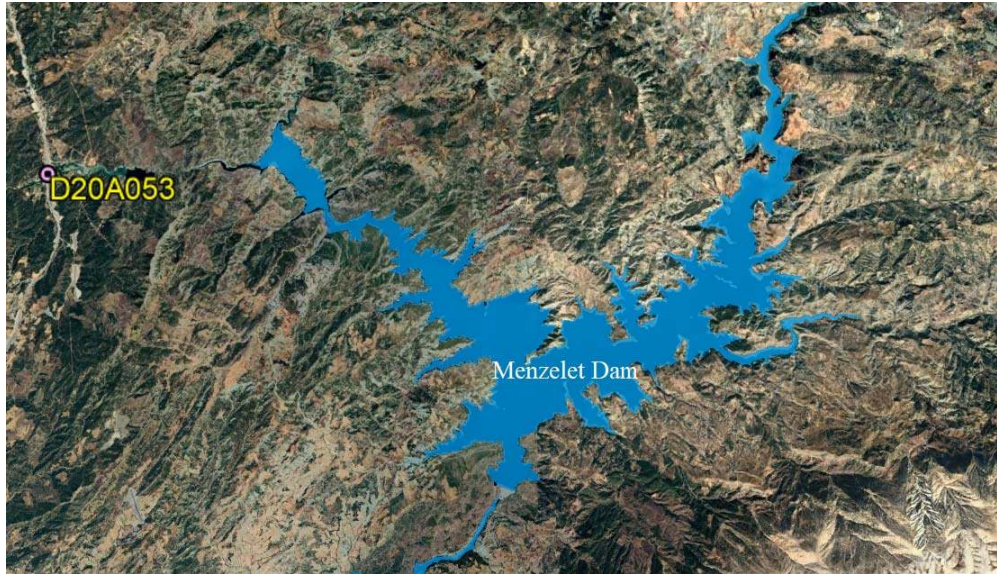


Figure 3.8 An illustration of location of Station and the nearest dam Menzelet

CHAPTER 4

RESULT AND DISCUSSION

4.1 Results for Imputation for Missing Data

Initially, computations were conducted using the monthly average discharge values from 1979 to 1994 as the training period and the monthly average discharge data from 1995 to 2014 as the test period. According to Figures 4.1 and 4.2, the R^2 values for the test period, which spans from 1995 to 2014, and the train period, which spans from 1979 to 1994, are 0.556 and 0.598, respectively, when computed with missing data.

After missing data was extracted, the computation was performed using CatBoost with the remaining monthly average discharge values. Without taking into account the missing years, Figure 4.3 illustrates how to use CatBoost as training to generate the data for all years between 1979 and 2014; the R^2 value is determined to be 0.903. Therefore, the decision has been made to use this technique for the imputation of missing data. The 1989 data and the 1997–2010 data are missing for the years 1979–2014. The data for the previously mentioned missing years will be estimated for this study.

According to Figures 4.4 and 4.5, calculation results after imputation of the R^2 values for the train period of 1979–1989 and 1997–2010 are 0.696, while the test period of 1990–1996 and 2011–2014 is determined to be 0.615. Table 4.1 illustrates the results of comparison of Root Mean Square Error (RMSE) and Mean Absolute Error (MAE) results of train and test periods before/after imputation and percentage differences. As can be seen from the table, there is a difference of 17.15% and 11.03%, respectively, in the RMSE and MAE values calculated for the test period. In other words, the imputation performed with CatBoost has shown a significant increase in accuracy.

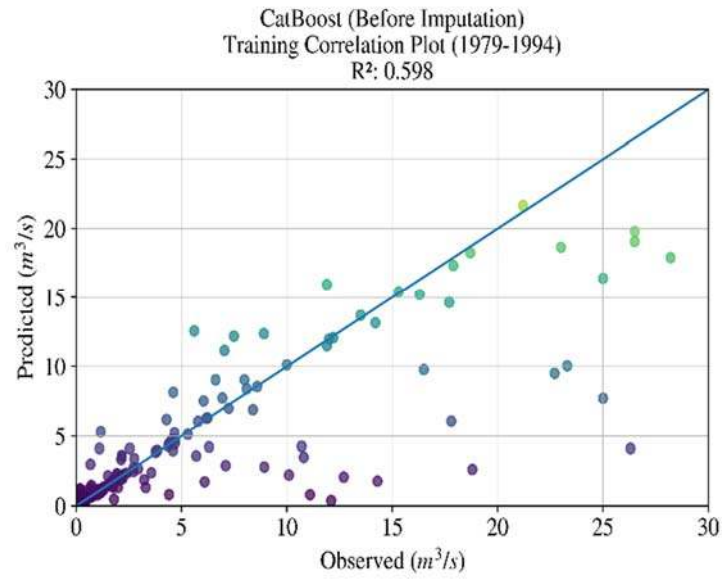


Figure 4.1 Scatterplot of train period 1979-1994 result before imputation

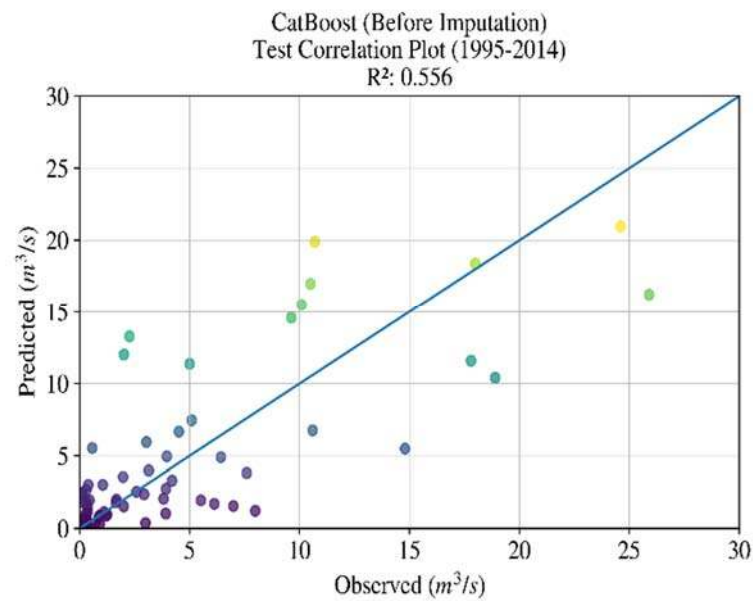


Figure 4.2 Scatterplot of test period 1995-2014 result after imputation

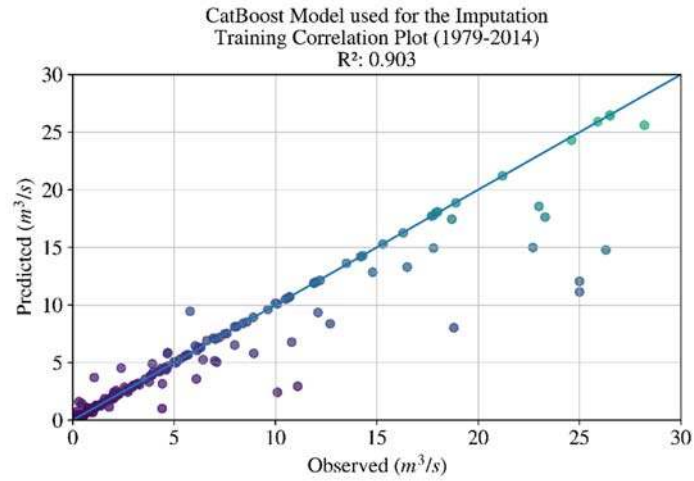


Figure 4.3 Scatterplot of train period 1979-2014 result before imputation

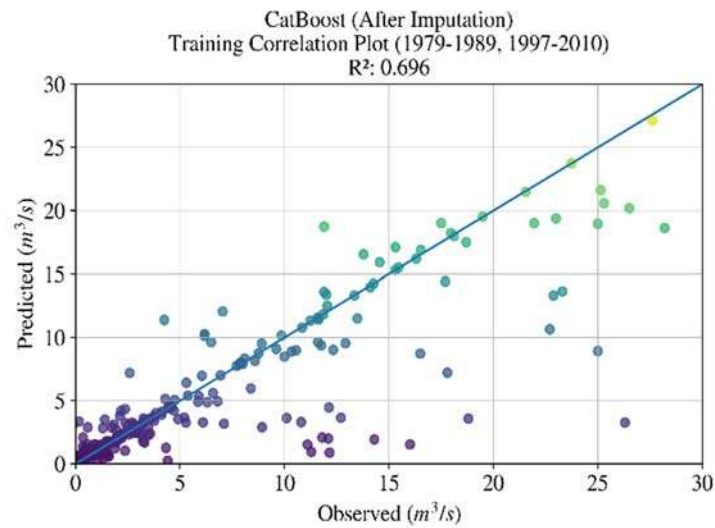


Figure 4.4 Scatterplot of train period result after imputation

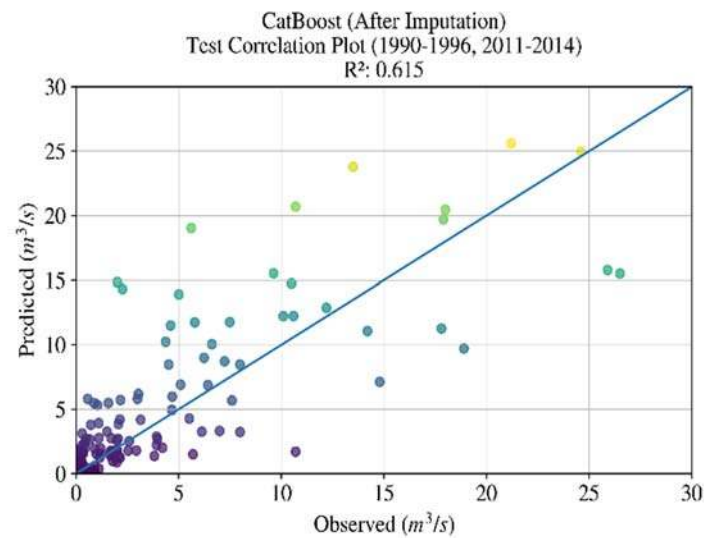


Figure 4.5 Scatterplot of test period result after imputation

Table 4.1 Comparison of RMSE and MAE results of train and test period before/after imputation

	Train Period			Test Period		
	Before Imputation	After Imputation	% Differences	Before Imputation	After Imputation	% Differences
RMSE	4.881	4.126	15.47	4.897	4.057	17.15
MAE	2.116	1.739	17.82	2.765	2.460	11.03

Figures 4.6 and 4.7 show the relationship between each of the 26 inputs and the mean discharge relationship both before and after imputation. As seen in Figure 4.6, which is before imputation, the inputs with the greatest weight in the mean discharge rate prediction are, in order: mean sea level pressure (mslp), 850 hPa relative vorticity of true wind (p8_z), 1000 hPa specific humidity (shum), 850 hPa specific humidity (s850), air temperature at 2 m (temp), while 850 hPa wind direction (p8th) is less. As seen in Figure 4.7 which is after imputation, the inputs with the greatest weight in the mean discharge rate prediction are, in order: 850 hPa relative vorticity of true wind (p8_z), mean sea level pressure (mslp), 1000 hPa specific humidity (shum), 850 hPa specific humidity (s850), 1000 hPa relative vorticity of true wind (p1_z), while 500 hPa wind direction (p5th) is less. As can be understood from this comparison, the contribution of each of the 26 inputs to the prediction varies with imputation.

The simplest technique obtains 5.274 m³/s average flow rate if the average flow for all other years is computed without taking the lost years into account. Figure 4.8 shows that how the correlation coefficient with the inputs would be if the calculations for the lost years were continued based on the average value of 5.274 m³/s using the simplest method. As seen in the figure, the inputs with the greatest weight in the mean discharge rate prediction are, in order: mean sea level pressure (mslp), 850 hPa relative vorticity of true wind (p8_z), 1000 hPa specific humidity (shum), 850 hPa specific humidity (s850), air temperature at 2 m (temp), while 850 hPa wind direction (p8th) is less. It means the weights of inputs would be like before the imputation figure if the mean discharges employed for missing years.

Figures 4.9 and 4.10 show the average impact results before and after imputation, respectively. Graphs depict the magnitude and direction of feature contributions across all predictions. The graphs show that the most crucial component is the mean sea level pressure (mslp) for before imputation while it is 500 hPa geopotential (p500) for after imputation.

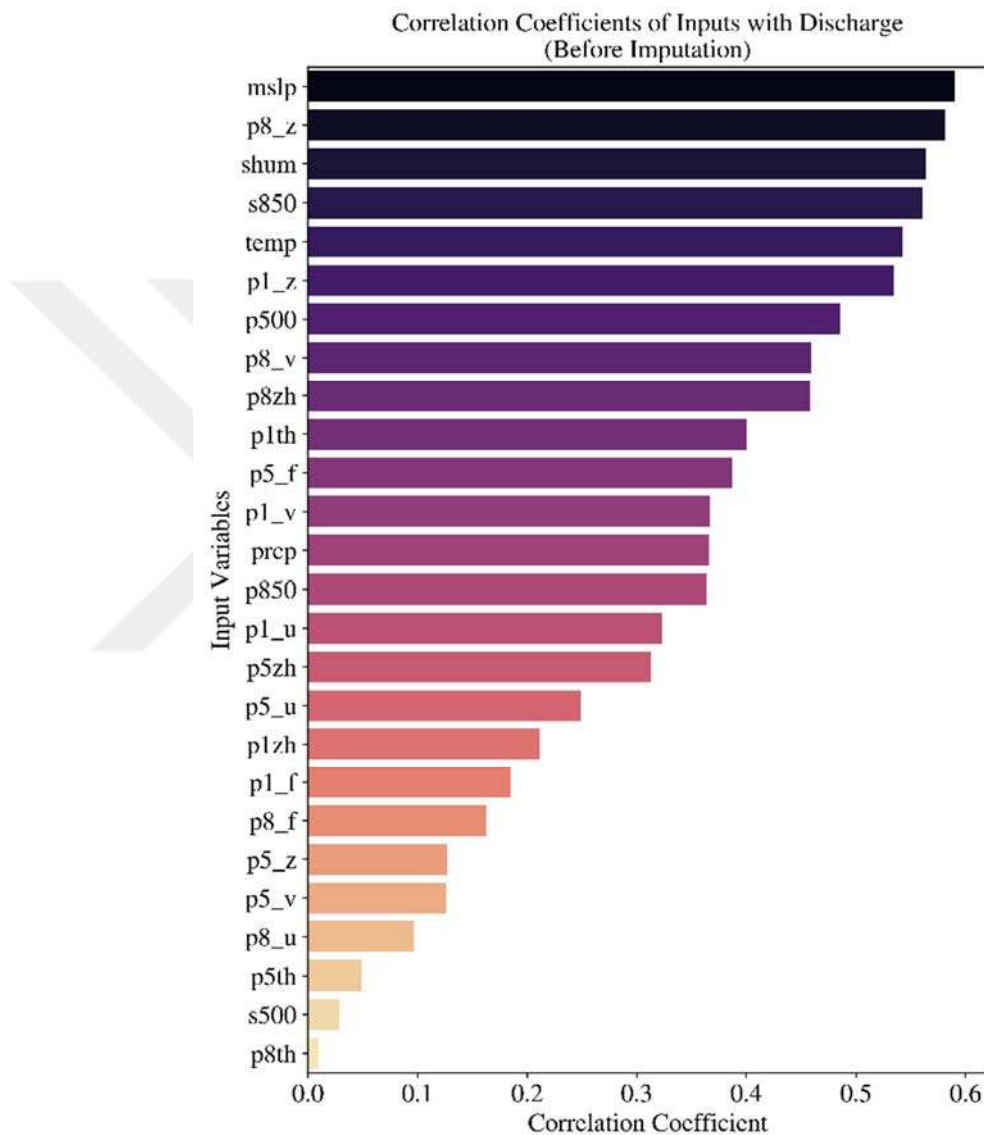


Figure 4.6 Correlation coefficient of inputs with discharge before imputation

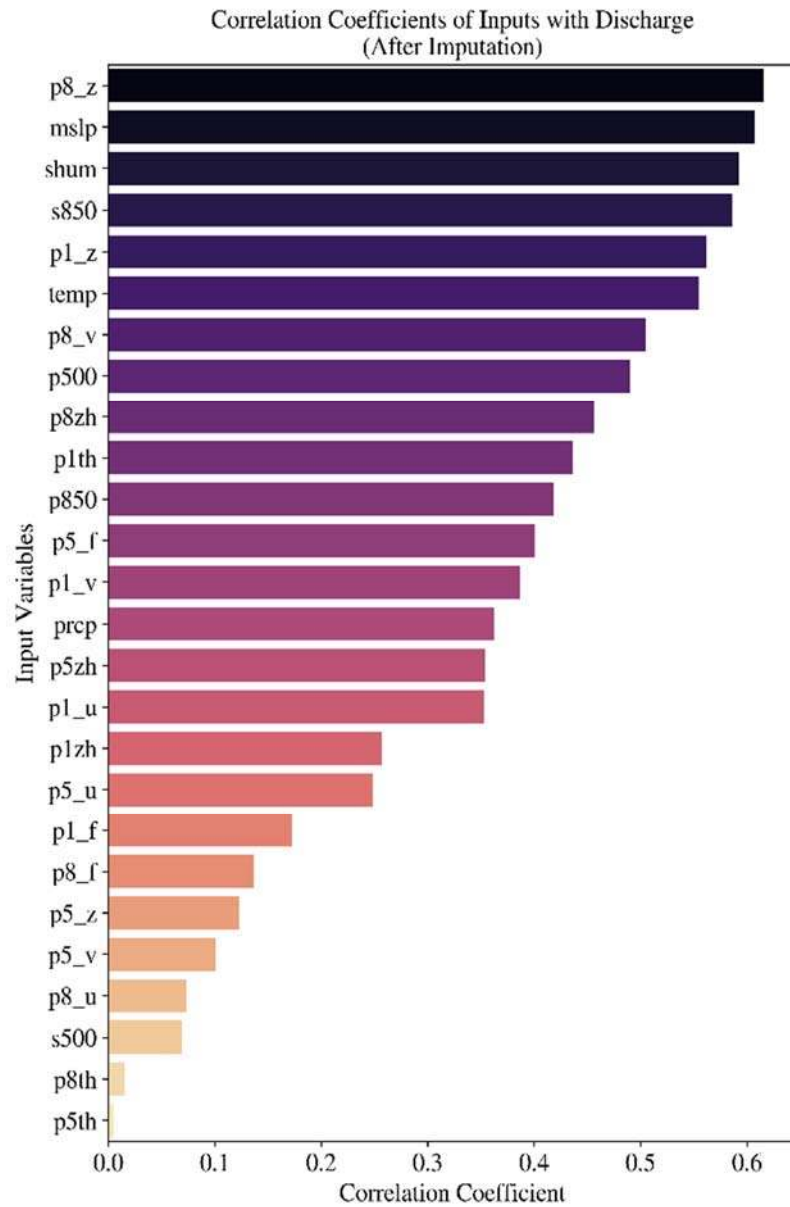


Figure 4.7 Correlation coefficients of inputs with discharge after imputation

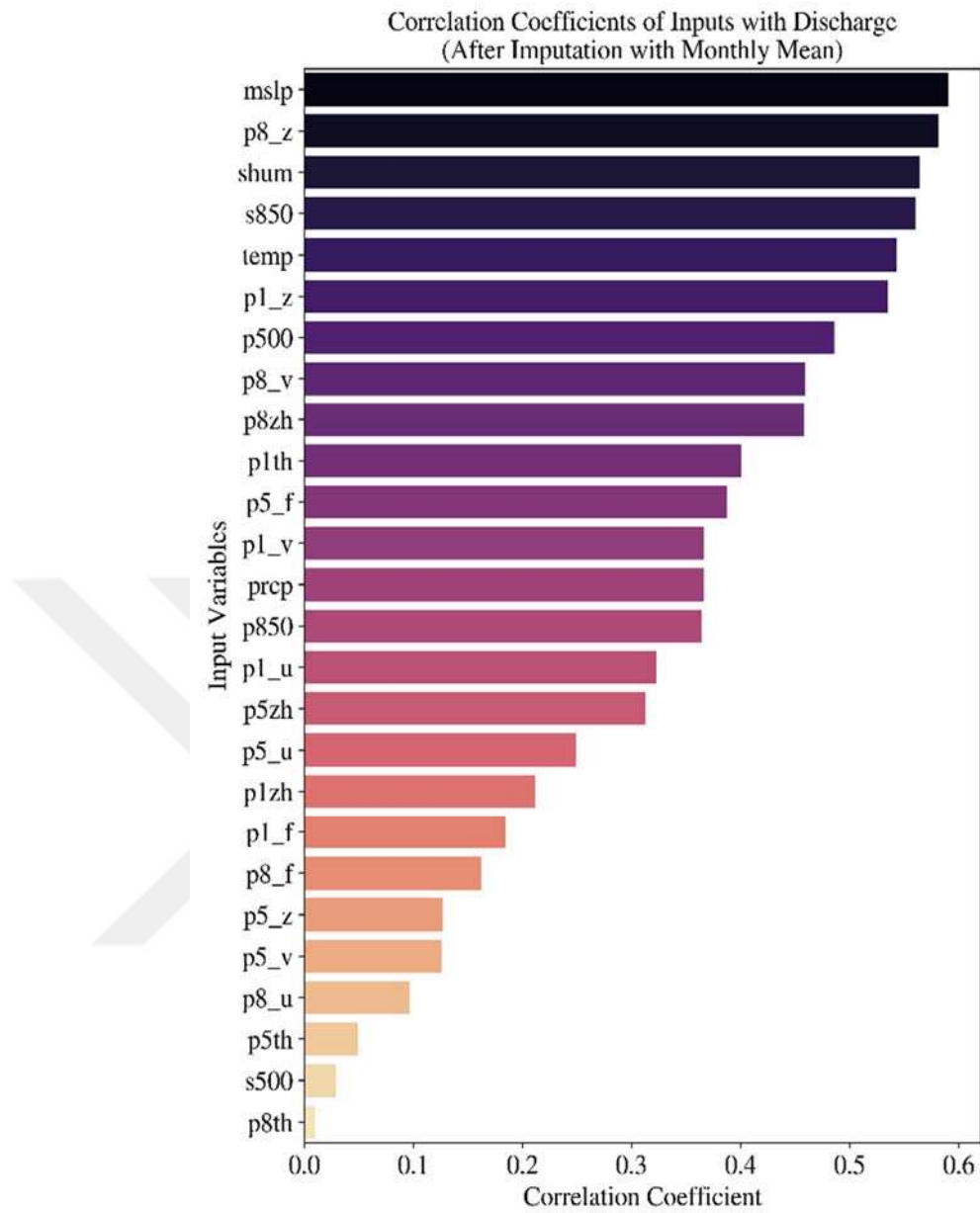


Figure 4.8 Correlation coefficients using simplest method after imputation with monthly mean discharge

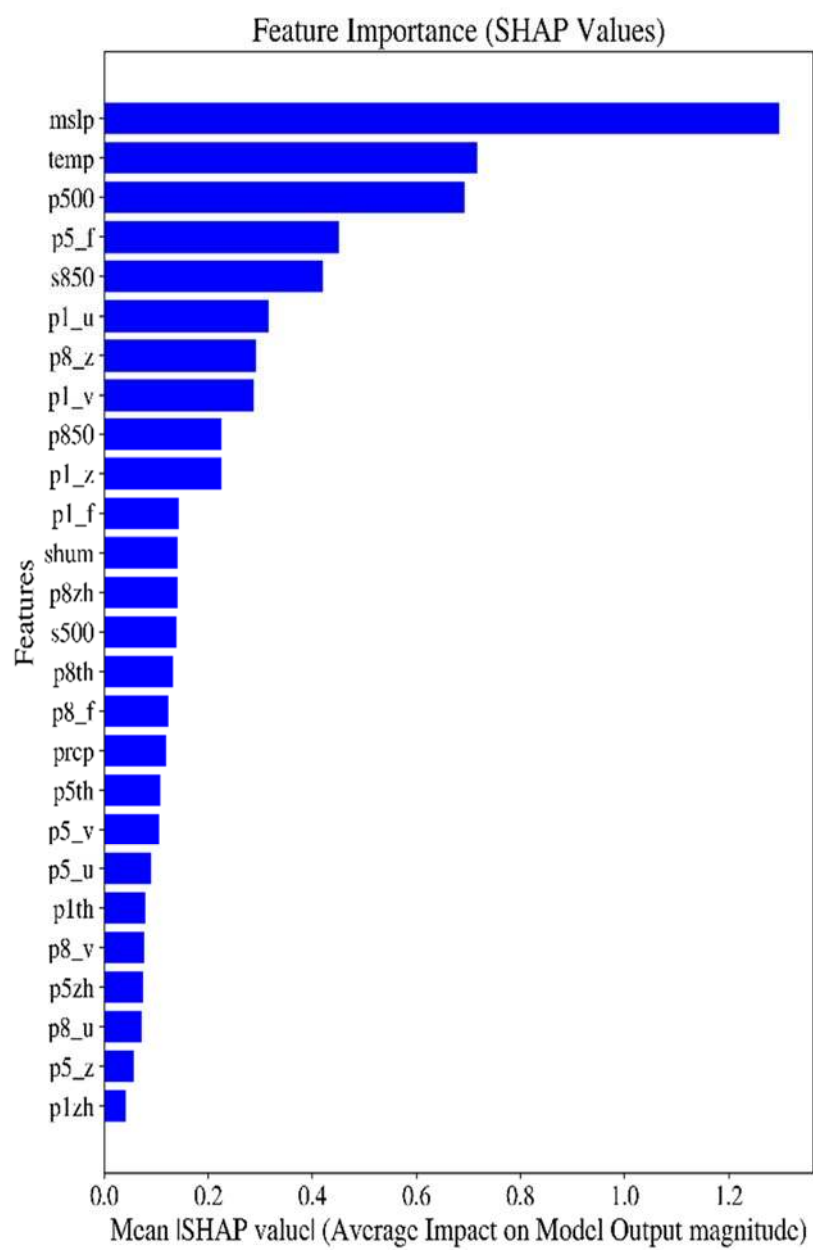


Figure 4.9 Average Impact Results Before Imputation

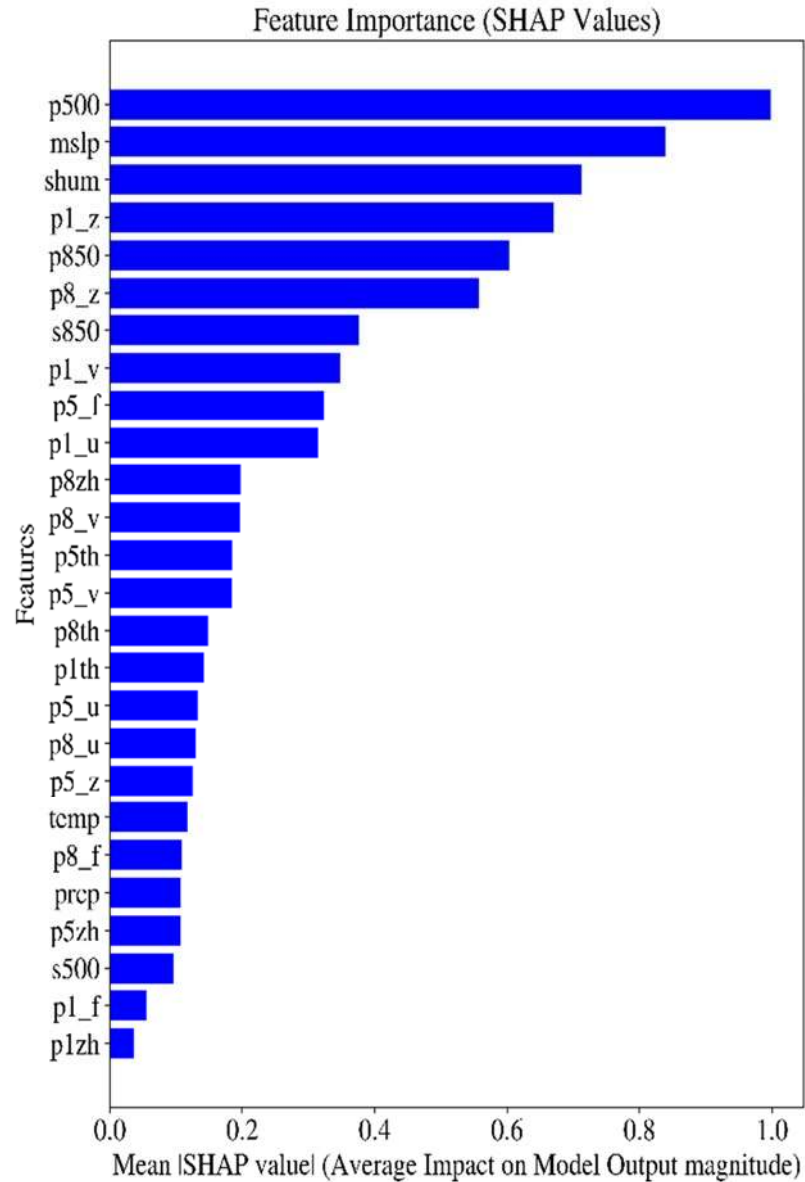


Figure 4.10 Average Impact Results After Imputation

A SHAP graph depicting how each factor effects a certain prediction. Figure 4.11 shows the SHAP graph before imputation, which depicts the magnitude and direction of feature contributions across all predictions. The color-coded scheme helps highlight the potential impact of each feature. Red represents higher feature values, while blue represents lower feature values. The graph shows that the most crucial component is the mean sea level pressure (mslp). As the graph shows that increasing the mean sea level pressure (mslp) value significantly reduces the discharge forecasts with 500 hPa Wind speed (p5_f) and 850 hPa Geopotential (p850) while the feature

air temperature at 2 m (temp), 500 hPa Geopotential (p500) and 850 hPa specific humidity (s850) a positive effect on the discharge predictions.

Figure 4.12 shows that after imputation, the most crucial component is the 500 hPa Geopotential (p500). As the graph shows that increasing value significantly affects the discharge forecasts positive side with 1000 hPa Relative vorticity of true wind (p1_z), 850 hPa Relative vorticity of true wind (p8_z), 850 hPa specific humidity (s850) while the feature Mean sea level pressure (mslp), 850 hPa Geopotential (p850) and 500 hPa Wind speed (p5_f) have negative effect on the discharge predictions.

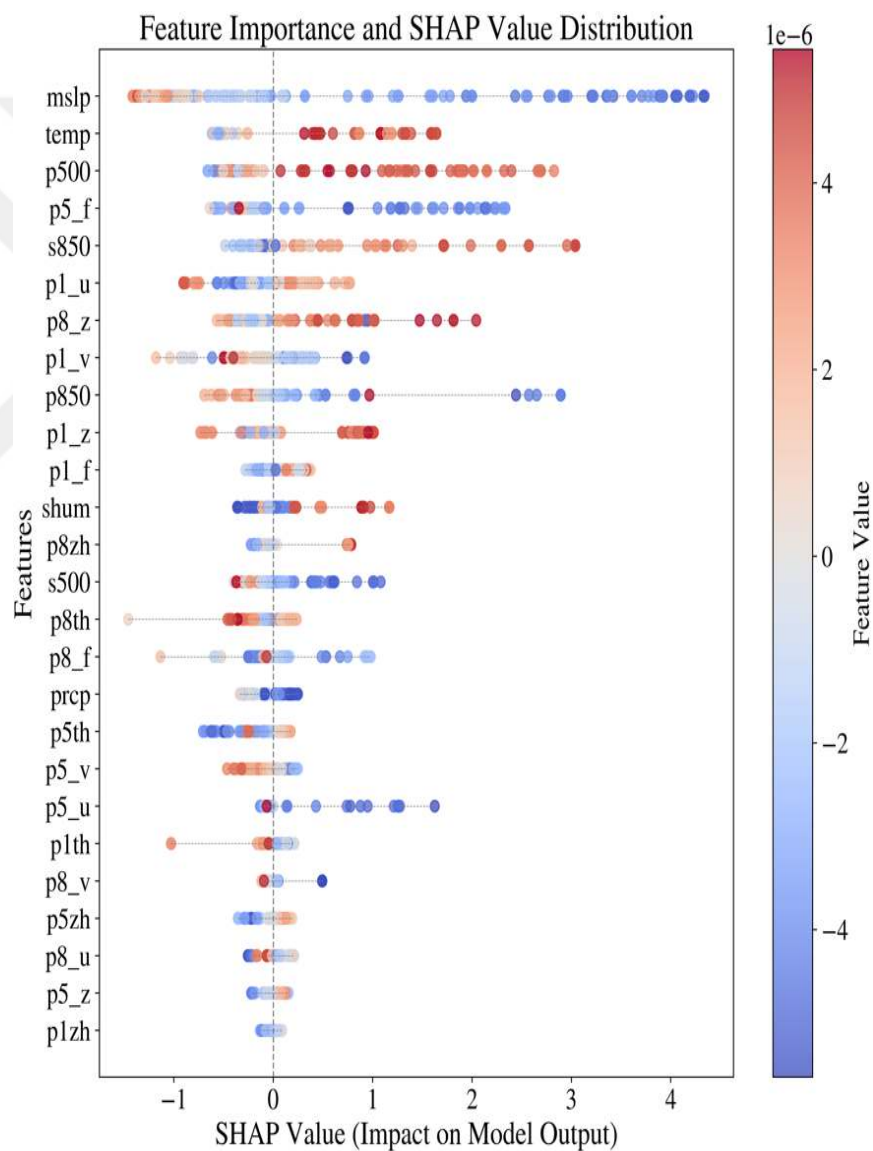


Figure 4.11 Impact results with SHAP value before imputation

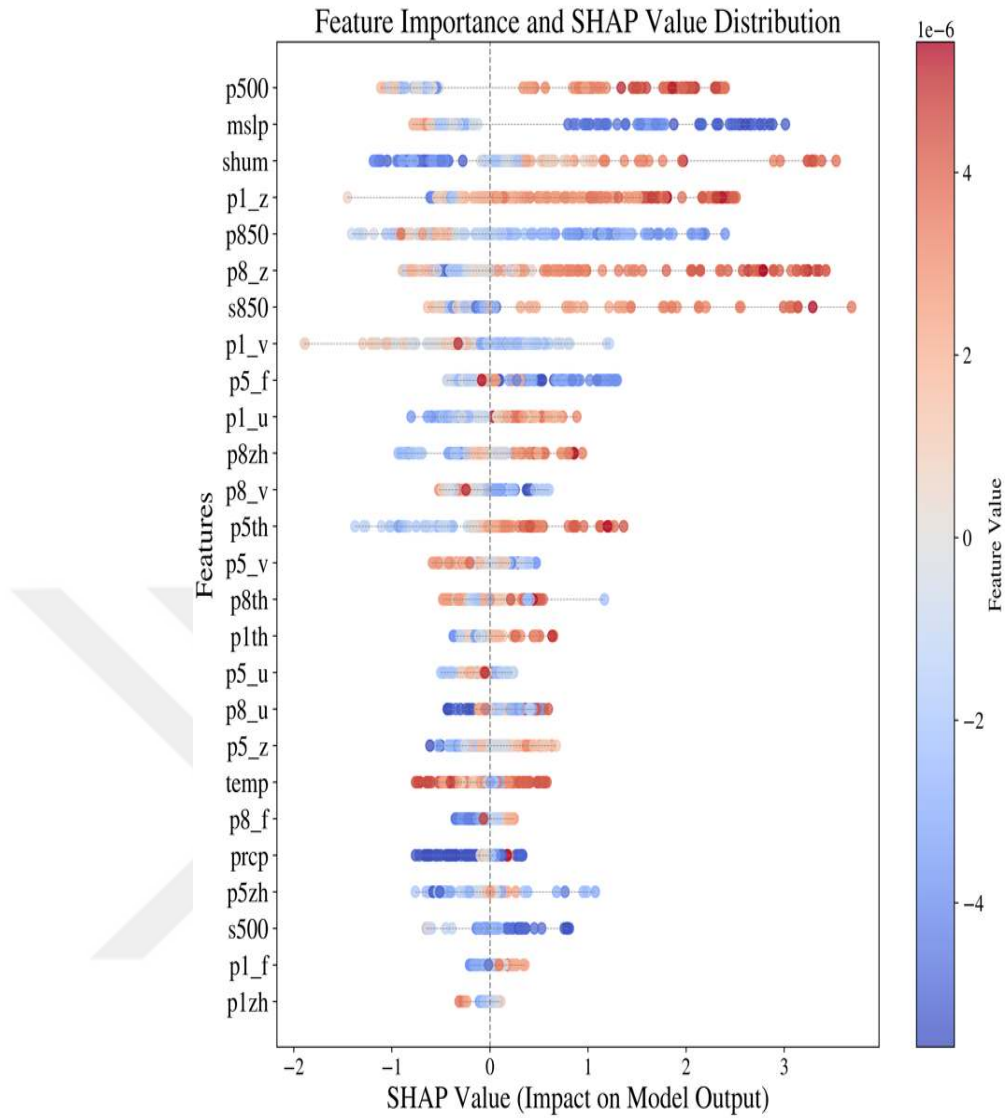


Figure 4.12 Impact results with SHAP value after imputation

4.2 Comparison of Methods Results

In this research, the large-scale climate variables (26 input sets) collected from the NorESM2-MM are utilized as the predictor for this method of estimation. Predictand uses local discharge data from the relevant station. Finally, a comparison between conventional models and CatBoost is performed. CatBoost was the best-performing model for discharge data during the testing periods. Table 4.5 compares the statistical results for the testing period as well as the observed values from 2007 to 2014. This comparison shows that GEP model provides the best result in predicting the mean and minimum values. The CatBoost model provides the best result in predicting the $NSE=0.626$. $RMSE=3.78 \text{ m}^3/\text{s}$ and $MAE=2.613 \text{ m}^3/\text{s}$ value, while Ridge Regression

model performs the poorest in predicting the all values. Furthermore, all models above the mean values and the best performance for predicting the standard deviation.

Figure 4.13 shows the relationship between each of the 26 inputs used during the projection and the mean discharge in the estimated scenario obtained. As seen in the figure, the inputs with the greatest weight in the mean discharge rate prediction are, in order: mean sea level pressure (mslp), air temperature at 2 m (temp), 1,000 hPa specific humidity (shum), 850 hPa specific humidity (s850), 850 hPa relative vorticity of true wind (p8_z), and 1,000 hPa relative vorticity of true wind (p1_z), respectively, while 500 hPa relative vorticity of true wind (p5_z) is less. As can be understood from this, the contribution of each of the 26 inputs to the prediction varies.

A SHAP graph depicting how each factor effects a certain prediction. Figure 4.14 shows the SHAP graph, which depicts the magnitude and direction of feature contributions across all predictions. The graph shows that the most crucial component is the mean sea level pressure (mslp). Also, Figure 4.15 shows that the direction of components.

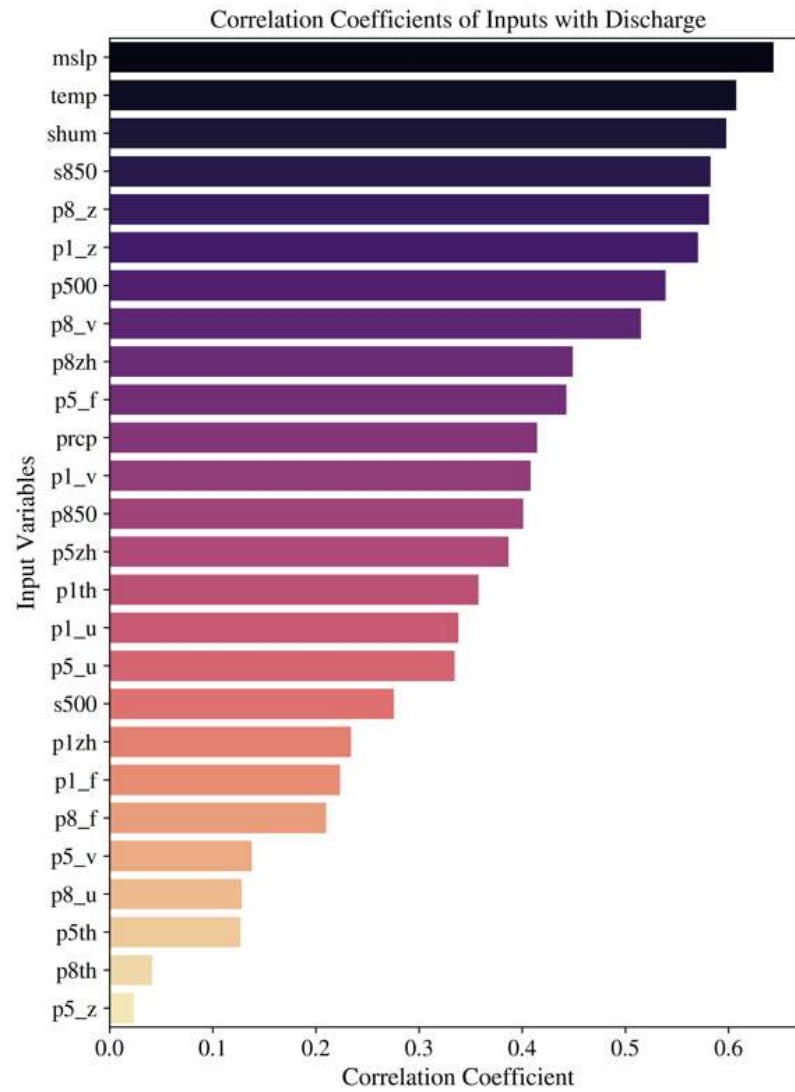


Figure 4.13 Correlation coefficient of inputs with discharge.

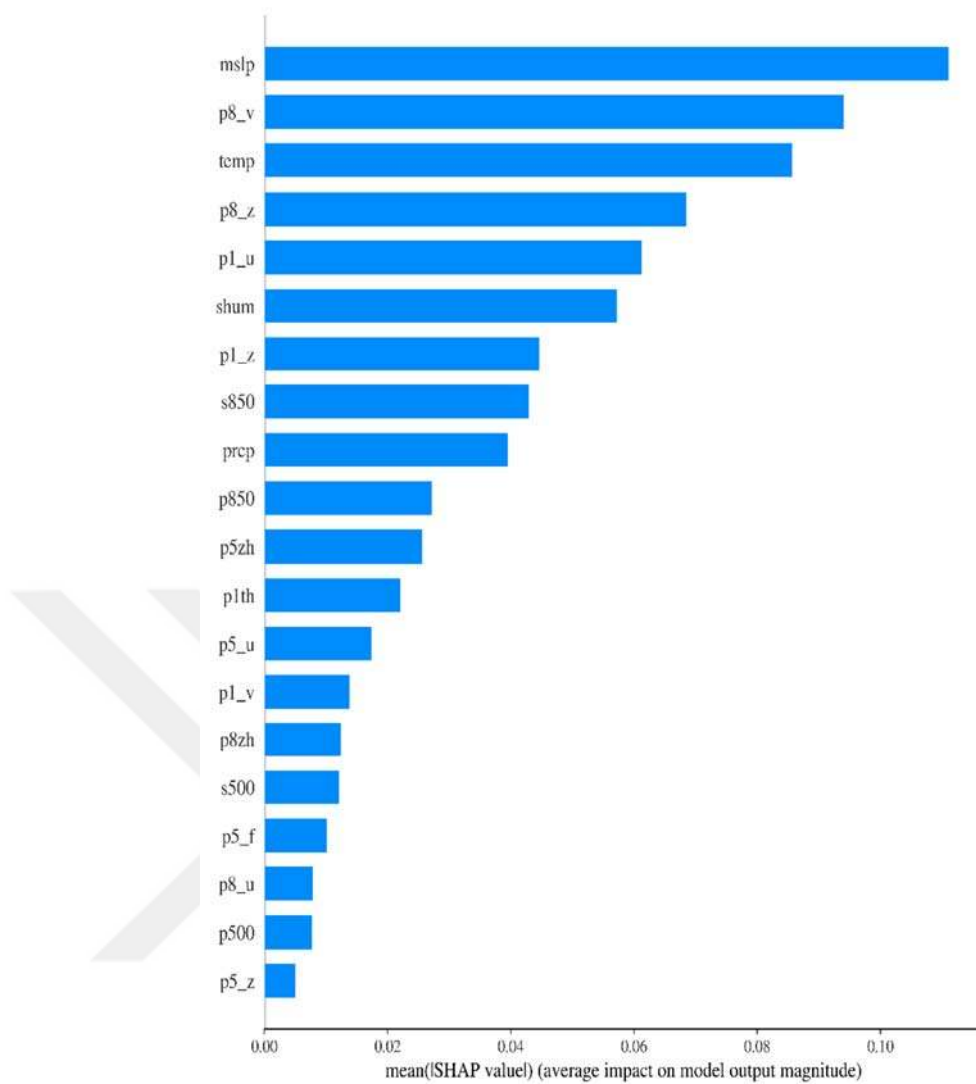


Figure 4.14 SHAP values average impact on model output magnitude.

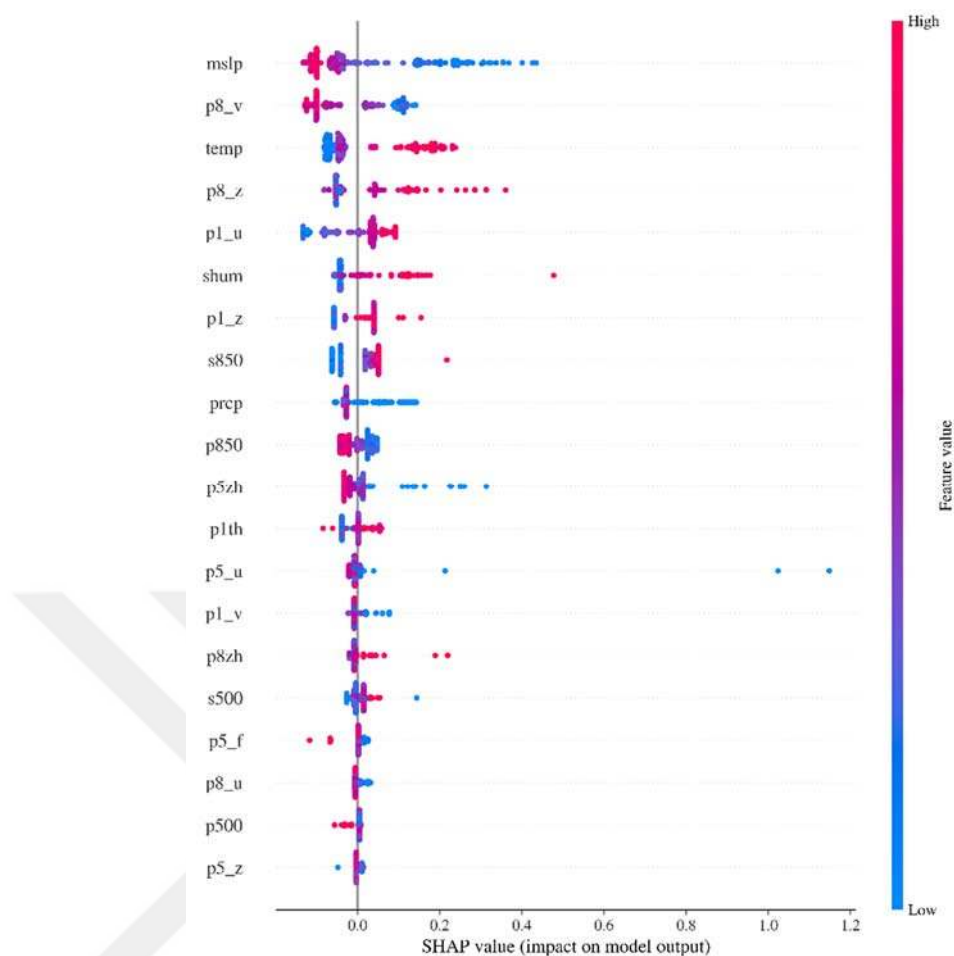


Figure 4.15 Impact results and direction of SHAP values

Table 4.2 Results of discharges for (m^3/s) for training period (1988-2006)

	Observed	GEP	RR	SVM	CatBoost
January	1.845	1.960747962	1.858842606	2.161270697	2.653273023
February	2.720555556	2.105108162	2.881196829	2.828523223	2.84404471
March	5.500555556	3.21228074	4.61297825	4.686852838	4.034761588
April	5.032777778	4.935113266	5.568122628	4.590462513	4.28863514
May	5.756111111	6.121813893	6.617327701	6.191050769	5.326856796
June	10.74333333	8.582174101	9.352429777	9.382634298	9.310540189
July	16.75944444	13.11257568	12.9264522	14.99722615	15.78011897
August	10.49833333	11.13198899	11.0231939	11.75372837	11.47487333
September	4.468333333	5.981757572	7.490381227	5.171654057	5.512369318
October	2.728333333	2.833073866	3.896552937	2.933695635	3.414018121
November	2.195555556	2.0441595	2.205271598	2.235302908	2.676740084
December	1.882777778	1.864751513	1.698361457	1.835830363	2.541979692

Table 4.3 Results of discharges(m^3/s) for testing period (2007-2014)

	Observed	GEP	RR	SVM	CatBoost
January	1.694286	1.97343773	2.1634807	1.5752623	2.56937811
February	2.24	2.34141067	3.1012561	3.3375916	3.21052593
March	3.988571	2.66320584	4.0687701	3.9743763	3.35137934
April	3.068571	4.47312179	5.7867824	5.055835	3.98763334
May	5.828571	6.761711	7.5477588	7.2340383	5.61287165
June	12.03571	9.42158185	9.8944359	9.6323673	9.74404106
July	16.36143	13.6982827	12.478539	13.943503	15.7309945
August	8.835714	11.7733957	11.282746	11.775372	12.8396612
September	3.418571	6.3393171	7.1600317	6.5740006	6.18293079
October	2.164286	2.86383424	3.761982	2.8623736	3.26218237
November	1.862857	1.97883443	2.0258322	1.8258329	2.5635185
December	1.641429	1.76546515	1.5241589	1.9939511	2.38035628

Table 4.4 Differences for testing period

	GEP	RR	SVM	CatBoost
	%differences	%differences	%differences	%differences
January	16.47609	27.69279	-7.02499	51.64964
February	4.527262	38.44893	48.99962	43.32705
March	-33.2291	2.010712	-0.3559	-15.9754
April	45.77213	88.5823	64.76185	29.95081
May	16.00975	29.49586	24.1134	-3.70073
June	-21.7198	-17.791	-19.9685	-19.0406
July	-16.277	-23.732	-14.7782	-3.85317
August	33.24781	27.69478	33.27017	45.31549
September	85.43761	109.4451	92.30257	80.863
October	32.32237	73.82095	32.25488	50.7279
November	6.225775	8.748662	-1.9875	37.61219
December	7.556623	-7.14437	21.47657	45.01735

Table 4.5 Comparison of model performance of observed and predicted Q (m³/s) for testing period between 2007-2014 years

Model	<i>KGE</i>	<i>RSR</i>	<i>NSE</i>	<i>RMSE</i>	<i>MAE</i>	Mean	Min	Max	Std. Deviation
Ridge Regression	0.453	0.775	0.399	4.788	3.117	5.9	1.524	12.479	3.771
GEP	0.517	0.748	0.44	4.621	2.779	5.504	1.765	13.698	4.146
SVM	0.545	0.734	0.461	4.534	2.957	5.815	1.575	13.944	4.113
CatBoost	0.650	0.611	0.626	3.78	2.613	5.953	2.380	15.731	4.459
Observed	-	-	-	-	-	5.262	1.641	16.361	4.749

Model performance results by month are demonstrated in Figure 4.16. The approach with the best test model performance was chosen, and using that approach, projections were created under several future climate change scenarios.

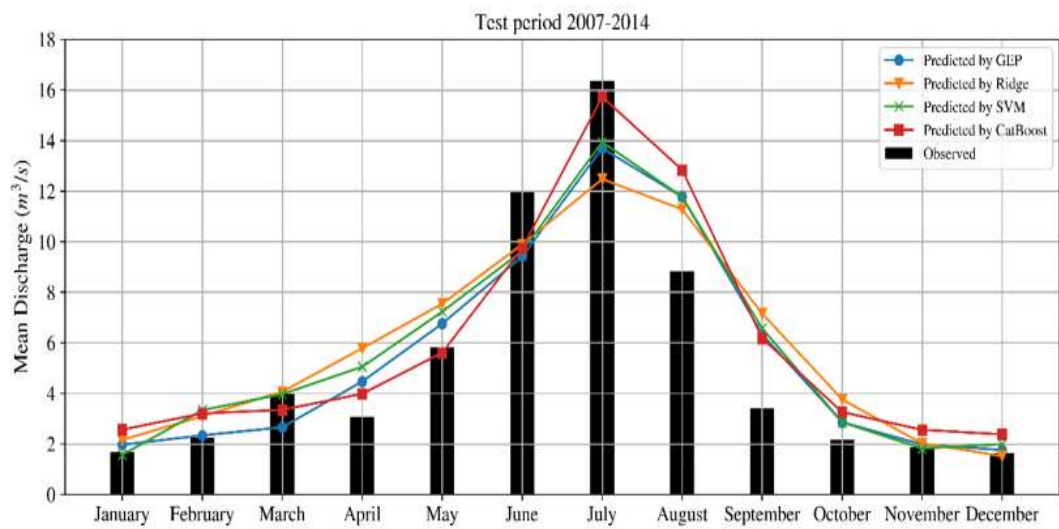


Figure 4.16 Observed and predicted monthly mean discharge performance for the testing period by all models

From Figure 4.17 to Figure 4.24 illustration of the scatter plots of observed and predicted mean discharge for GEP, Ridge Regression, SVM, and CatBoost for the training and testing period are given, respectively. The CatBoost model performed the best with the highest NSE value of 0.626, RMSE 3.78 m³/s, and MAE 2.613 m³/s values, while the ridge regression model performs the poorest in predicting all values. Considering that the KGE value shows good performance as it approaches 1 and the

RSR value shows good performance as it approaches 1, it is observed that CatBoost again provided the best result with KGE 0.650 and RSR 0.611. In predicting the maximum and standard deviation values, all models made predictions below the observed value.

Also, from Figure 4.25 to Figure 4.32 illustration of line graphs of comparison of the observed and predicted discharges for GEP, Ridge Regression, SVM, and CatBoost for the training and testing period are given, respectively.

The CatBoost model provided the closest predictions except for the minimum value. The CatBoost model was superior to the others in estimating NSE, RMSE, and MAE statistical values.

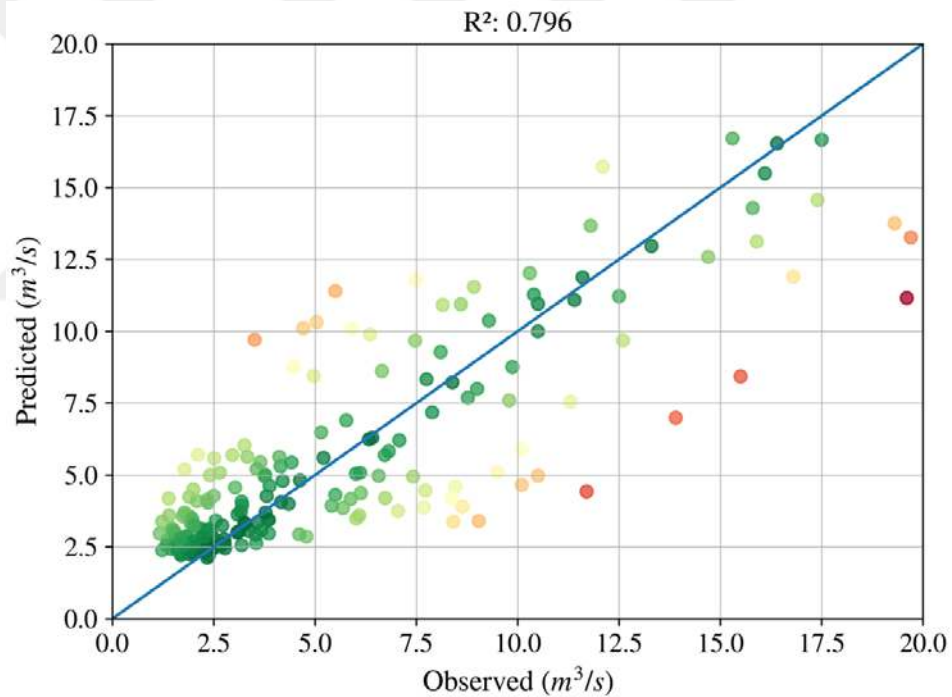


Figure 4.17 Scatterplot of observed and predicted Q for training period with CatBoost

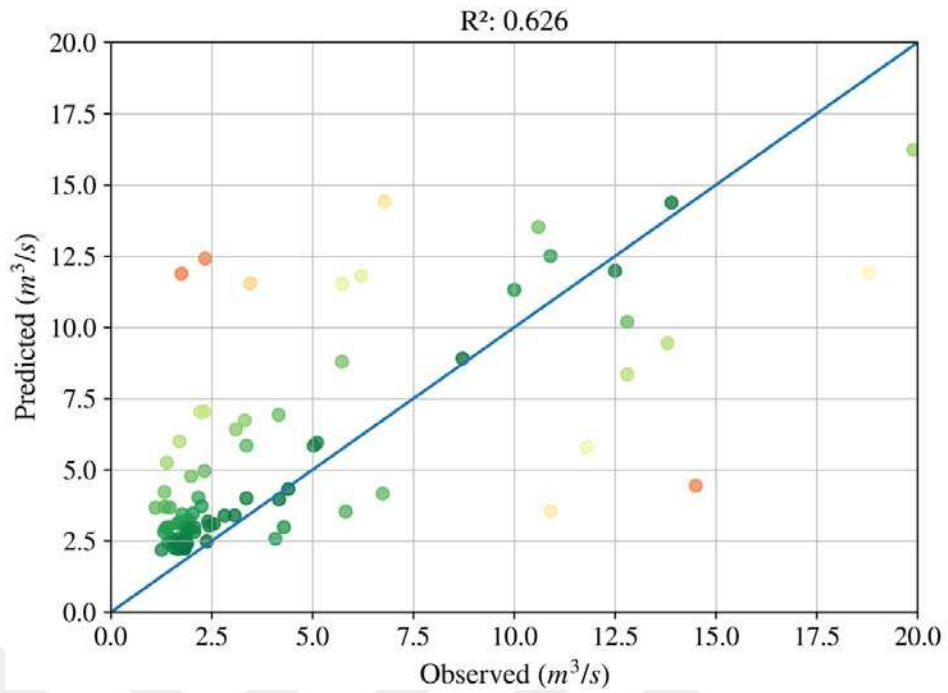


Figure 4.18 Scatterplot of observed and predicted Q for testing period with CatBoost

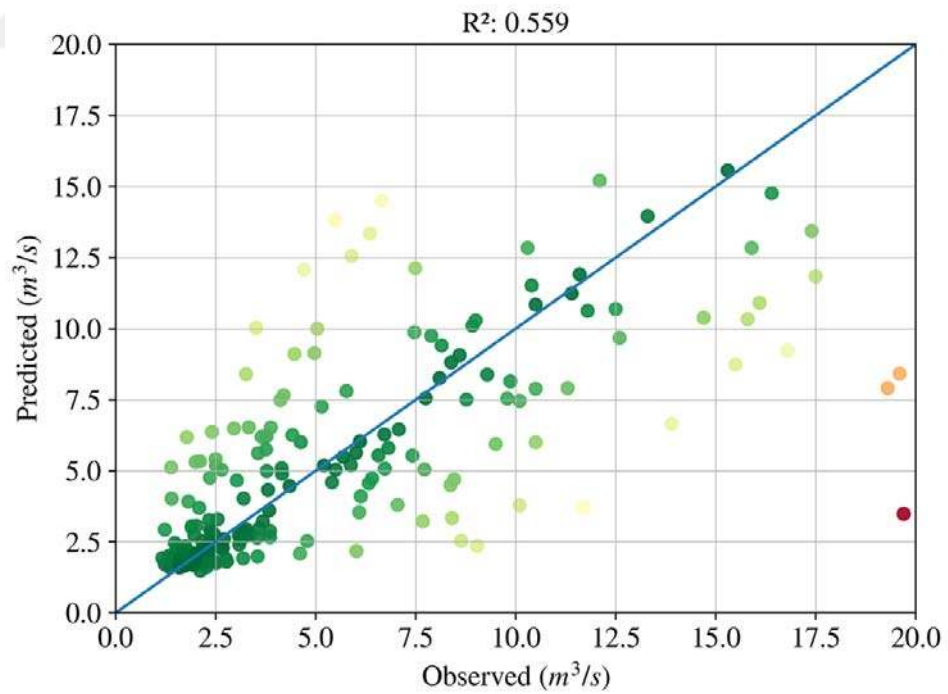


Figure 4.19 Scatterplot of observed and predicted Q for training period with GEP

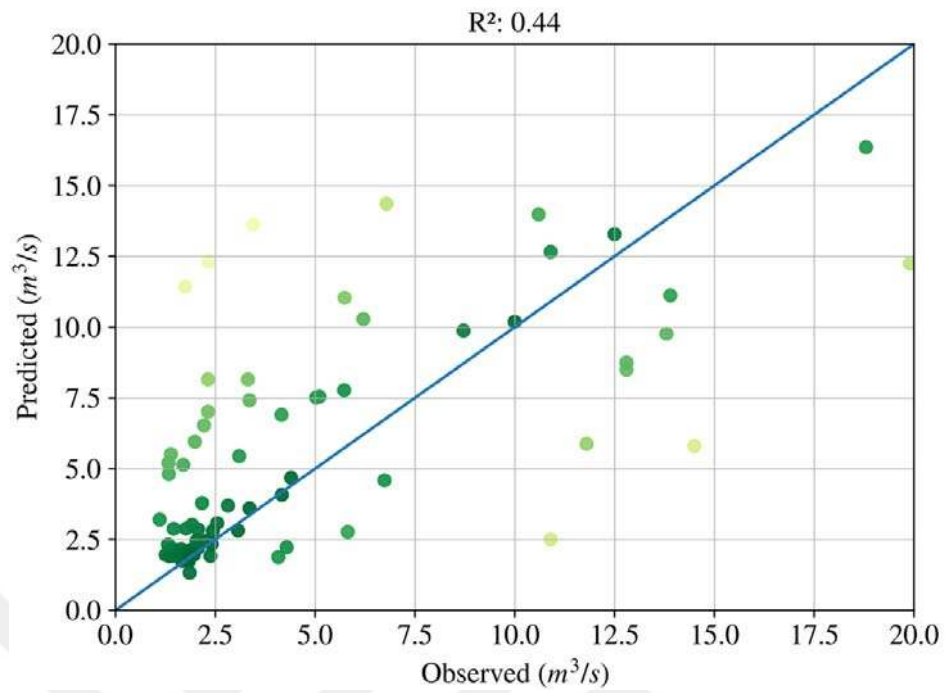


Figure 4.20 Scatterplot of observed and predicted Q for testing period with GEP

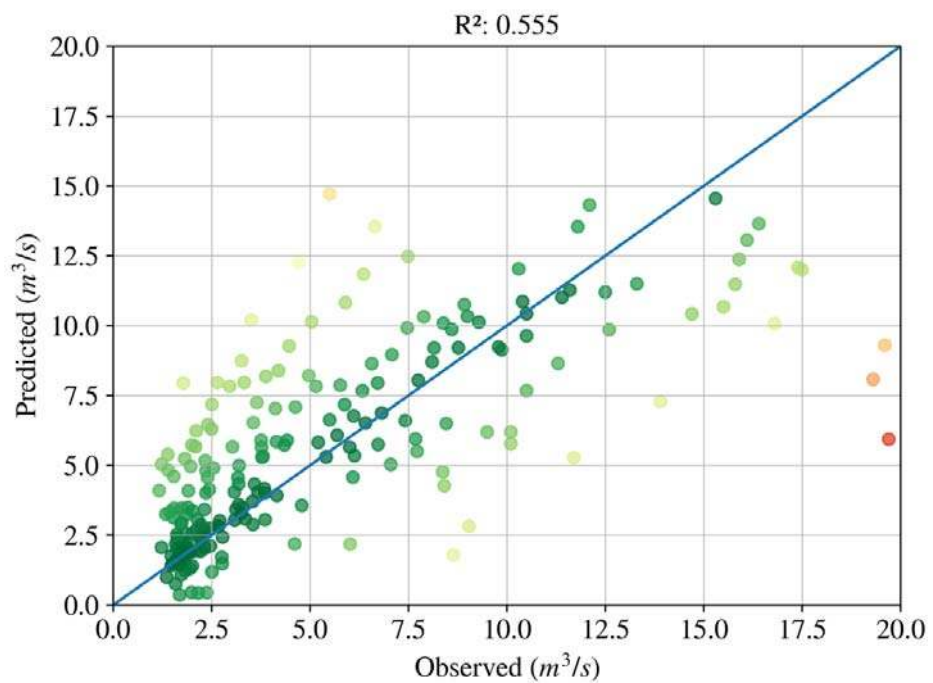


Figure 4.21 Scatterplot of observed and predicted Q for training period with Ridge Regression

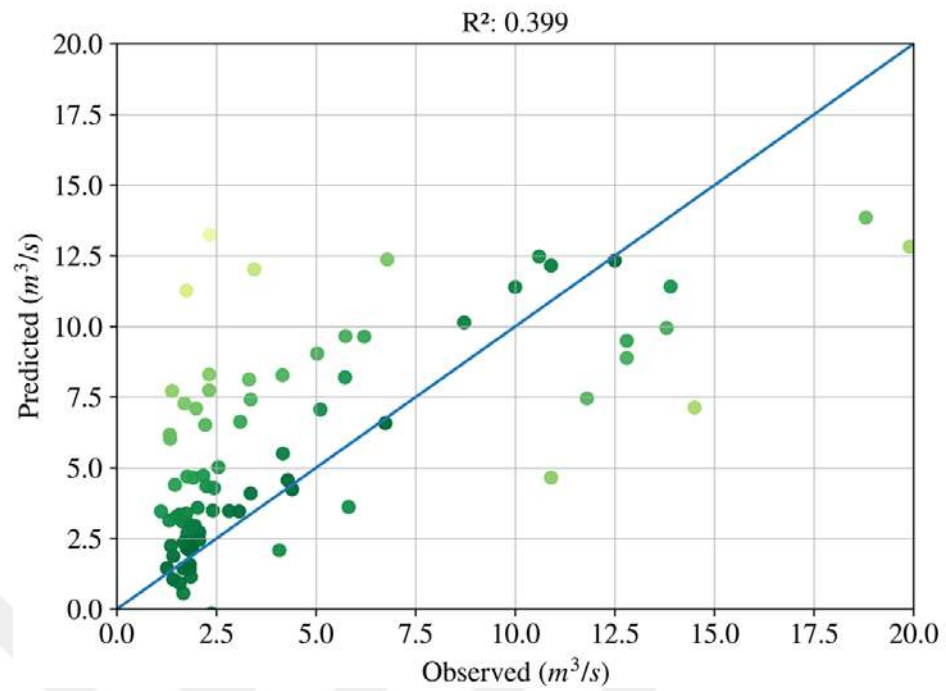


Figure 4.22 Scatterplot of observed and predicted Q for testing period with Ridge Regression

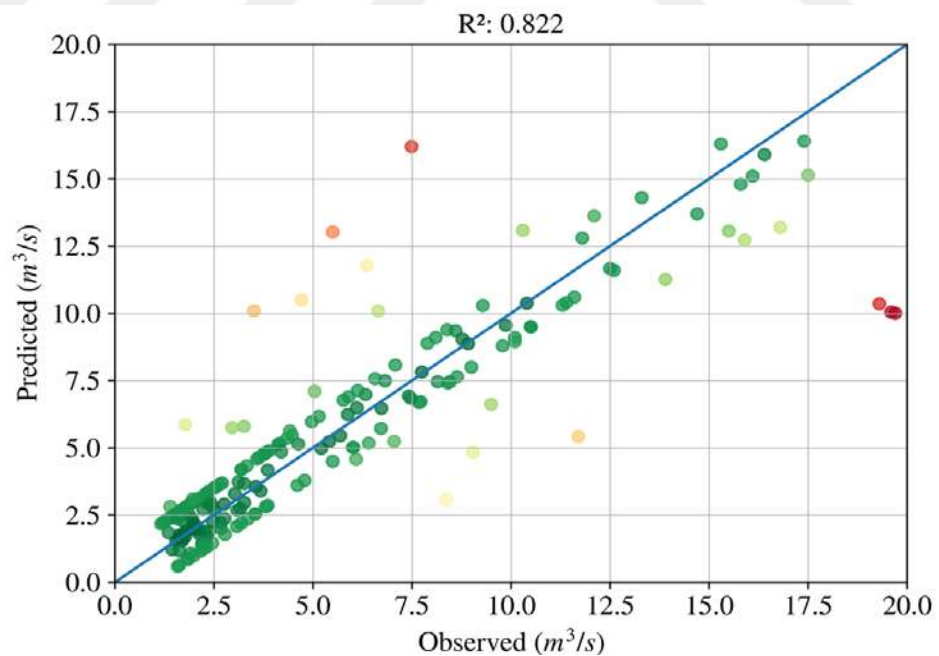


Figure 4.23 Scatterplot of observed and predicted Q for training period with SVM

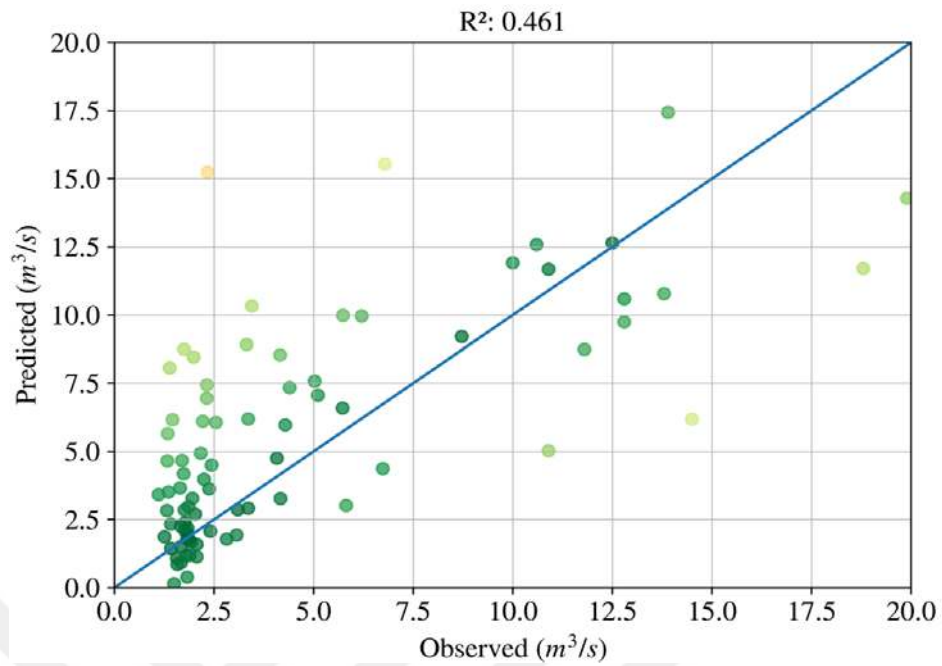


Figure 4.24 Scatterplot of observed and predicted Q for testing period with SVM

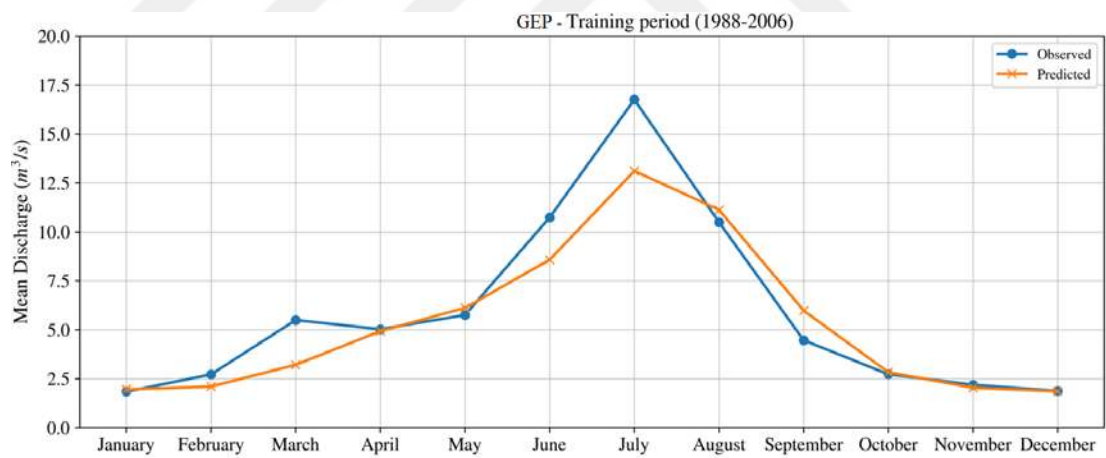


Figure 4.25 Comparison of the result of observed and predicted discharges for training period in GEP method

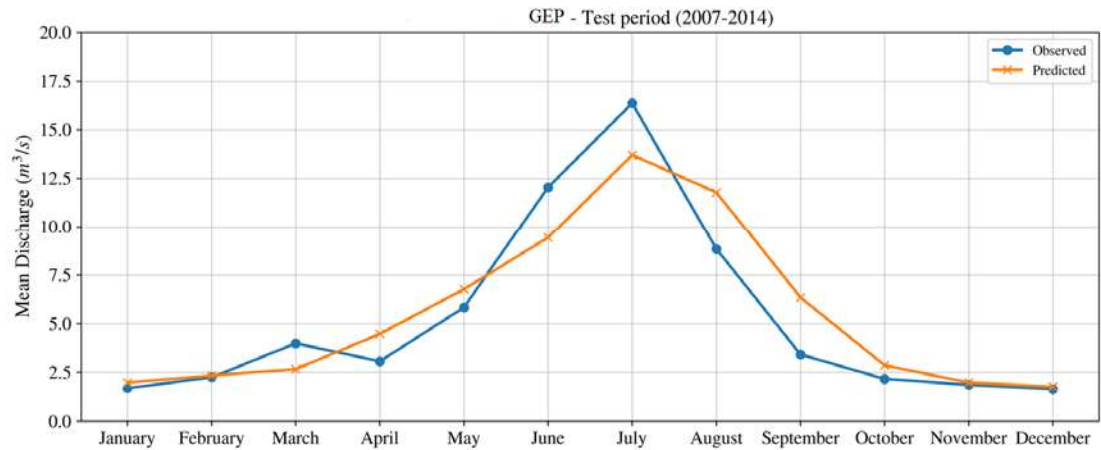


Figure 4.26 Comparison of the result of observed and predicted discharges for test period in GEP method

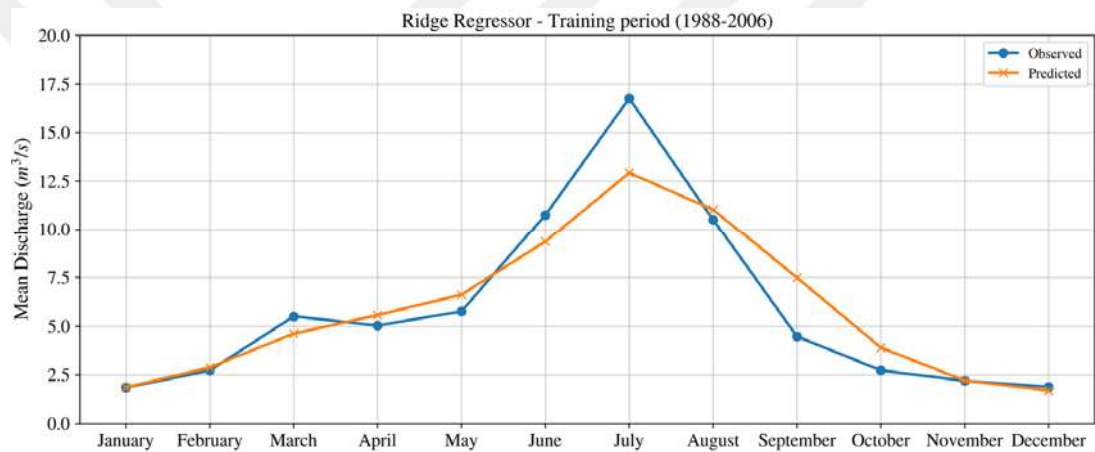


Figure 4.27 Comparison of the result of observed and predicted discharges for training period in Ridge Regression method

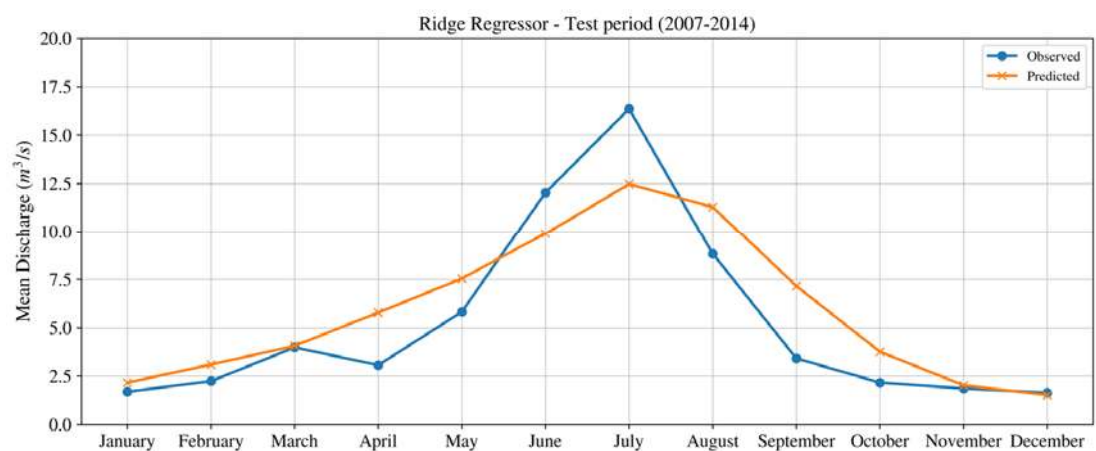


Figure 4.28 Comparison of the result of observed and predicted discharges for test period in Ridge Regression method

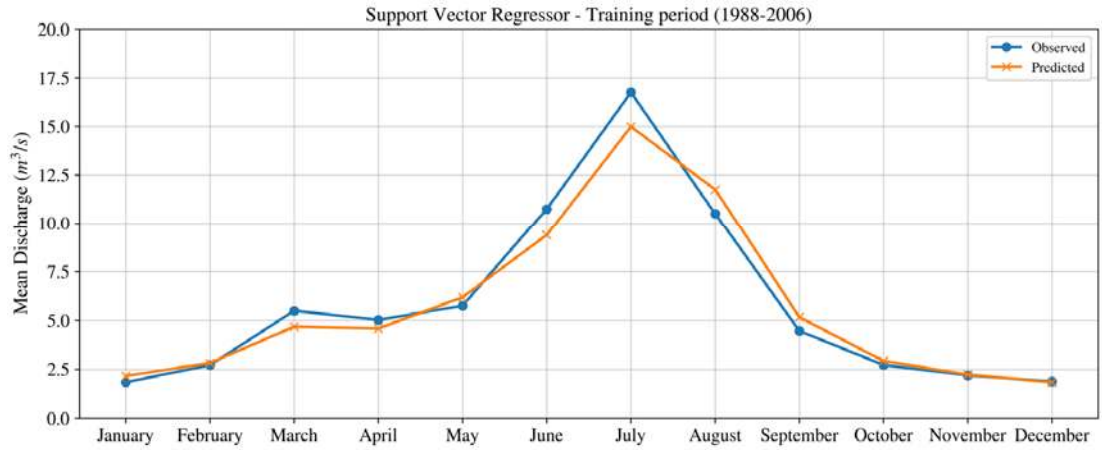


Figure 4.29 Comparison of the result of observed and predicted discharges for training period in SVM method

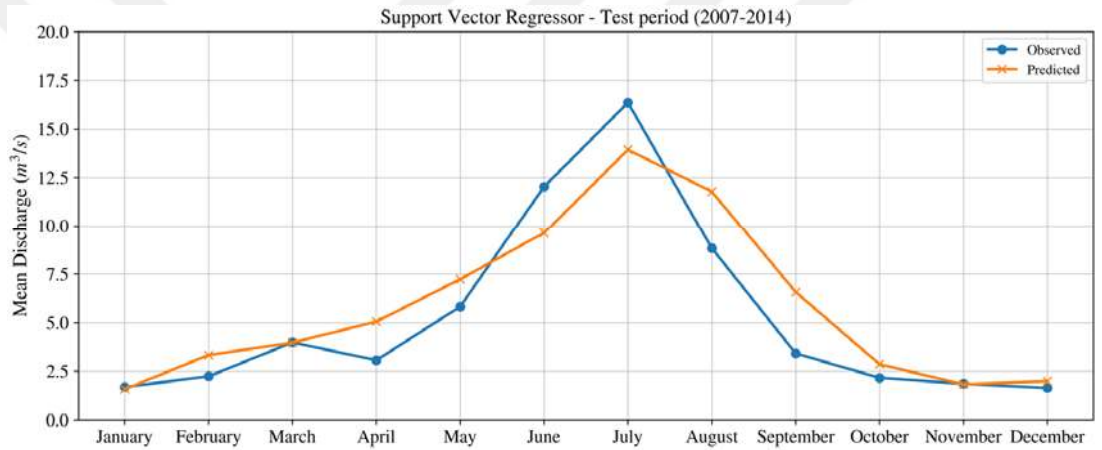


Figure 4.30 Comparison of the result of observed and predicted discharges for test period in SVM method

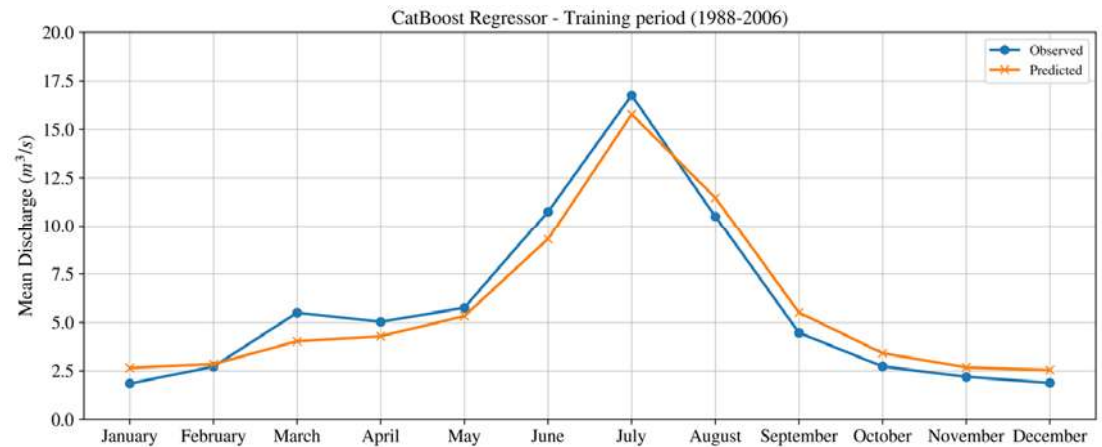


Figure 4.31 Comparison of the result of observed and predicted discharges for training period in CatBoost method

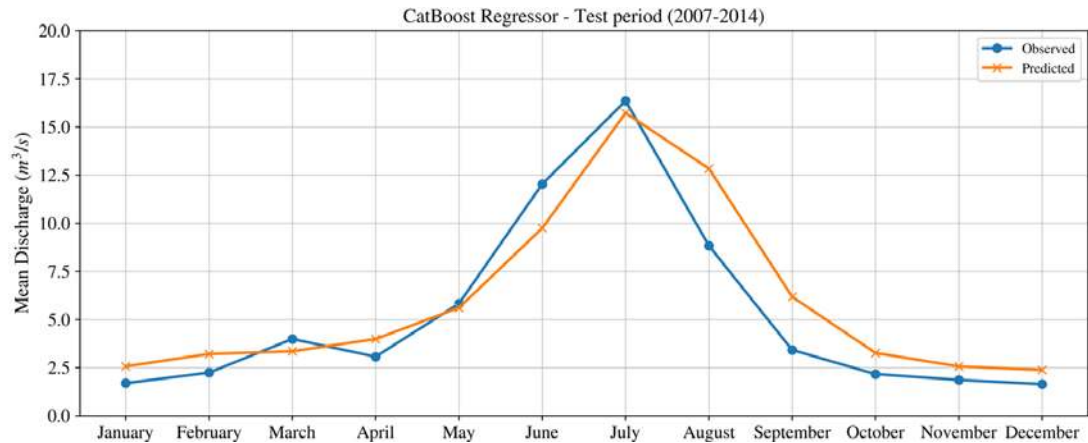


Figure 4.32 Comparison of the result of observed and predicted discharges for test period in CatBoost method

In comparison to the observed data, the discharge values were typically overestimated by the models, as shown in Figure 4.16. While all models understated discharge in June and July, they all overestimated it in February, April, August, September, and October. Other months—January, March, May, November, and December—have variations aside from these months. The SVM model was undervalued in January, the Ridge Regression was overvalued in March, and in May, the opposite was true. SVM, Ridge Regression, and CatBoost were all undervalued in November, December, and November, respectively. Overestimation percentages for the GEP model's performance ranged from 85.43% in September to 4.52% in February, which was the best performance.

In March and November, the SVM model produced estimates that were 0.35 and 1.98% below the observed values, and 92.3% above the worst values in September. The Ridge Regression model performed the worst in September, outperforming the observed value by 109.44%. The CatBoost model produced extremely accurate estimates that were 3.70 and 3.85% below the observed values in May and July, respectively. All things considered, the CatBoost model outperformed the Ridge Regression model in terms of forecasting the observed values, whereas the latter demonstrated comparatively poor performance.

The chosen station's historical data was separated into train and test sets, and computations were performed using both a novel technique called CatBoost and more conventional models like GEP, SVM, and Ridge Regression. Compared to the

other approaches, the analyses showed that CatBoost produced more reliable results. Because of this, CatBoost has been used to carry out the study's primary goal, which is to project the future under climate change.

4.3 Future projection of discharge with CatBoost model under SSP scenarios

According to O'Neill et al. (2016), the SSP scenarios show different pathways for socioeconomic growth and greenhouse gas emissions that affect the state of the climate in the future. Shared Socioeconomic Pathways (SSPs), also known as scenarios, highlight particular situations and are usually associated with projections regarding socioeconomic trends, climate impacts, and greenhouse gas emissions. Below are the different types of scenarios and their contents.

- SSP1-2.6: Outlines a scenario with low emissions and aggressive climate goals aiming to raise the temperature by roughly 2°C over pre-industrial levels.
- SSP2-4.5: A scenario with medium emissions and a moderate climate policy that raises temperatures by 2.5–3°C.
- SSP3-7.0: Indicates a high-emissions scenario with significant greenhouse gas emissions and little climate policy. possibly leading to a significant warming of 3.5–4 degrees Celsius.
- SSP5-8.5, which emphasizes the use of fossil fuels and has minimal climate regulations, is an exceptionally high emissions scenario that will cause more warming. possibly going above 4°C.

Comprehensive assessment frameworks use these scenarios to forecast possible climate effects based on a variety of hypotheses regarding upcoming socioeconomic shifts and policy choices.

NorESM2-MM in the context of CMIP6. making certain that they are able to offer trustworthy climate projections for the future that are supported by solid statistical evidence. Clearly defining these time periods improves the model's legitimacy and helps guide decisions about adaptation and climate policy. Three time periods—2015–2039, 2040–2069, and 2070–2100—have been created from the projected data between 2015 and 2100.

First, the observed data for each SSP has been compared with the established priorities. and then the projected scenarios for each period have been compared within themselves. From Figure 4.34 to Figure 4.40 the line graphs illustrate a comparison of the monthly forecast results for 2015 – 2039, 2040 – 2069, 2070 - 2100 periods with the observed mean discharge m^3/s between 1988 and 2014 years and the predicted data for under SSP1-2.6, SSP2-4.5, SSP3-7.0, and SSP5-8.5, respectively. Each line graph includes baseline observed data between 1988 and 2014 years. For all line graphs, mean discharges have reached the highest rainfall in July; a sharp decline was observed from July to September. and a gradual decrease has been noted from September to December. For all line graphs. the common point is that mean discharge for all periods has increased in January, February, and between August and December while decreasing between March and July.

Also, Table 4.6, Table 4.7, and 4.8 show predicted discharges values for each month under SSP scenarios between the years 2015-2039, 2040-2069 and 2070-2100, respectively. Table 4.9 demonstrates that the comparison of mean and maximum predicted discharges values under SSP scenarios between 2015-2100 years.

Table 4.6 Comparison of predicted discharges (m^3/s) between 2015-2039 years under each scenario with CatBoost

	SSP126	SSP245	SSP370	SSP585
January	2.718606	2.989465	2.907864271	2.643
February	2.956364	2.901947	3.131714283	3.0532
March	3.496736	3.462672	3.507731078	3.594
April	4.330275	4.324446	4.352496495	4.0987
May	5.399764	5.227848	5.508528301	5.165
June	10.30408	10.0614	9.987854364	10.662
July	14.34809	14.25218	13.8467165	13.291
August	12.26157	12.6086	12.50563086	12.719
September	6.383179	6.144855	6.599059277	6.3239
October	3.473154	3.501524	3.655370628	3.433
November	2.910469	2.83662	2.73681745	2.7534
December	2.658957	2.683399	2.608451684	2.6374

Table 4.7 Comparison of predicted discharges (m³/s) between 2040-2069 years under each scenario with CatBoost

	SSP126	SSP245	SSP370	SSP585
January	2.587866	2.63159	2.60474	2.880604
February	2.91818	2.87256	2.853556	2.8071
March	3.408679	3.46725	3.654383	3.482076
April	4.047362	4.05681	4.697864	4.258196
May	6.041832	5.61703	5.626318	5.526635
June	10.6074	10.4725	10.97309	10.60477
July	15.32261	15.0162	15.26093	14.47085
August	12.32059	12.9201	12.76079	12.72667
September	6.54851	6.90535	7.259409	7.19796
October	3.485281	3.67716	3.557933	3.691469
November	2.873955	2.81668	2.794904	2.854515
December	2.554701	2.60815	2.607484	2.671033

Table 4.8 Comparison of predicted discharges (m³/s) between 2070-2100 years under each scenario with CatBoost

	SSP126	SSP245	SSP370	SSP585
January	2.643185	2.899334	2.887092	2.858962
February	2.70566	3.020545	3.024829	2.885278
March	3.382276	3.362039	3.615288	3.81462
April	4.659697	4.505705	4.653675	4.462054
May	5.2264	5.986807	6.023414	5.864255
June	10.04686	10.41255	10.86855	11.37048
July	14.84379	13.01614	17.00471	16.75727
August	12.29502	12.35396	13.87211	14.5454
September	6.908454	7.570819	7.934721	7.593448
October	3.527737	3.680995	3.864079	4.023229
November	2.78567	2.746885	3.087165	3.100059
December	2.600401	2.744527	2.689958	2.746575

Table 4.9 Comparison of future projection under different SSPs with observed Q (m^3/s) between 2015-2100 years.

	Observed		SSP1-2.6		SSP2-4.5		SSP3-7.0		SSP5-8.5	
Years	Mean	Max	Mean	Max	Mean	Max	Mean	Max	Mean	Max
2015-2039			5.936	14.348	5.916	14.252	5.945	13.846	5.864	13.291
2040-2069			6.059	15.322	6.088	15.260	6.221	15.261	6.097	14.470
2070-2100			5.968	14.843	6.025	13.016	6.627	17.004	6.668	16.757
1988-2014	5.681	32.8								
Overall	5.681	32.8	5.988	14.838	6.010	14.176	6.264	15.370	6.210	14.839

The monthly average flows for the SSP1-2.6 scenario are projected for three different time periods, including 2015 and 2100, as shown in Figure 4.33. It is known that the decline is less severe than what was seen in July and August. The forecast indicates a slightly different mean discharge for the 2040–2069 period ($6.059 \text{ m}^3/\text{s}$), compared to $5.968 \text{ m}^3/\text{s}$ for the 2070–2100 period and $5.936 \text{ m}^3/\text{s}$ for the 2015–2039 period. July $15.322 \text{ m}^3/\text{s}$ will be the highest mean discharge amount during the 2040–2069 time frame.

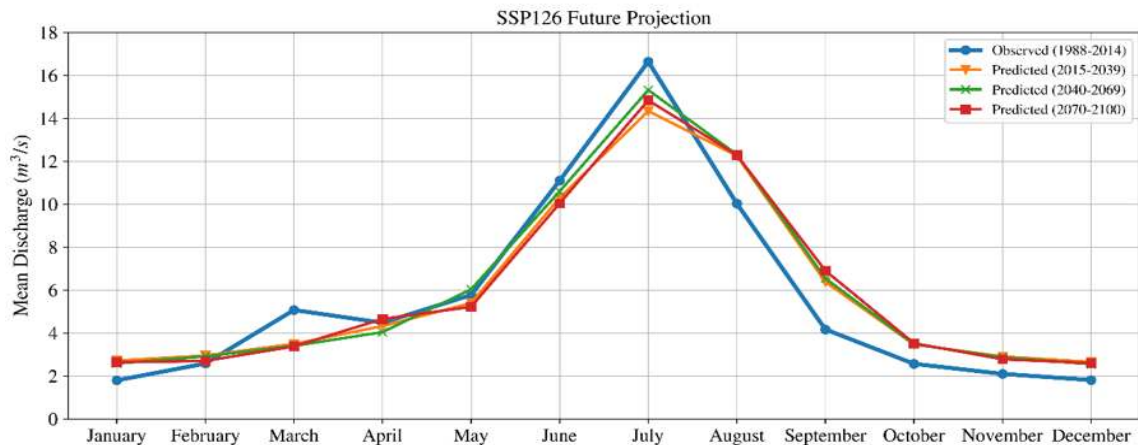


Figure 4.33 Line graph of observed and predicted mean discharge with default model parameters under SSP1-2.6 scenario

It is noted that the estimated data computed for the SSP2-4.5 scenario for the same periods, with the exception of March, parallels the observed data, and Figure 4.34 exhibits similarities with Figure 4.33. The decline is less severe than what was seen

in July and August. In a similar vein, it is evident that the 2040 and 2069 periods deviate slightly from the others and forecast higher mean discharges of $6.088 \text{ m}^3/\text{s}$, $6.025 \text{ m}^3/\text{s}$, and $5.916 \text{ m}^3/\text{s}$, respectively, for the 2070–2100 period and 2015–2039. July $15.260 \text{ m}^3/\text{s}$ will be the highest mean discharge amount during the 2040–2069 timeframe.

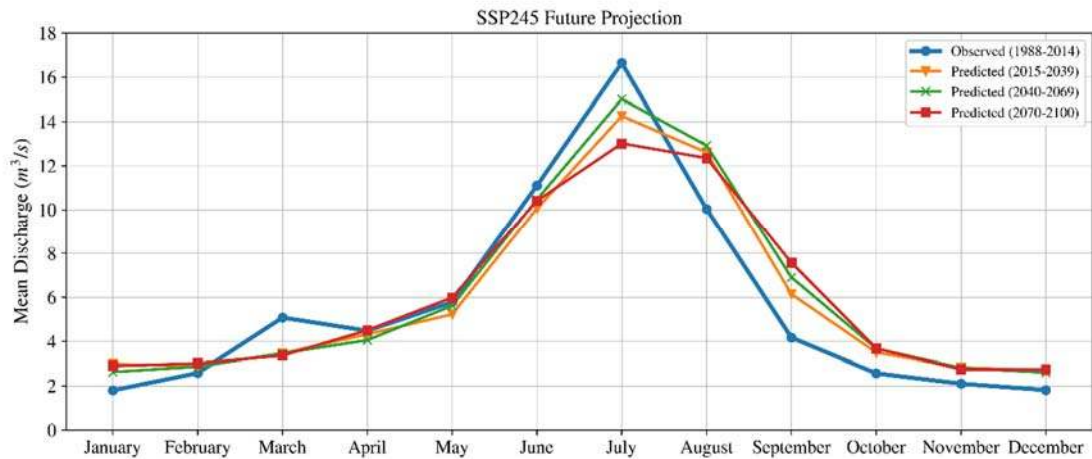


Figure 4.34 Line graph of observed and predicted mean discharge with default model parameters under SSP2-4.5 scenario

The expected mean discharge of SSP3-7.0 from 2015 to 2100 years is shown in Figure 4.35. While the 2015–2039 period predicted a mean discharge amount of $5.945 \text{ m}^3/\text{s}$, the 2070 and 2100 periods predicted $6.627 \text{ m}^3/\text{s}$, which was higher than previous periods.

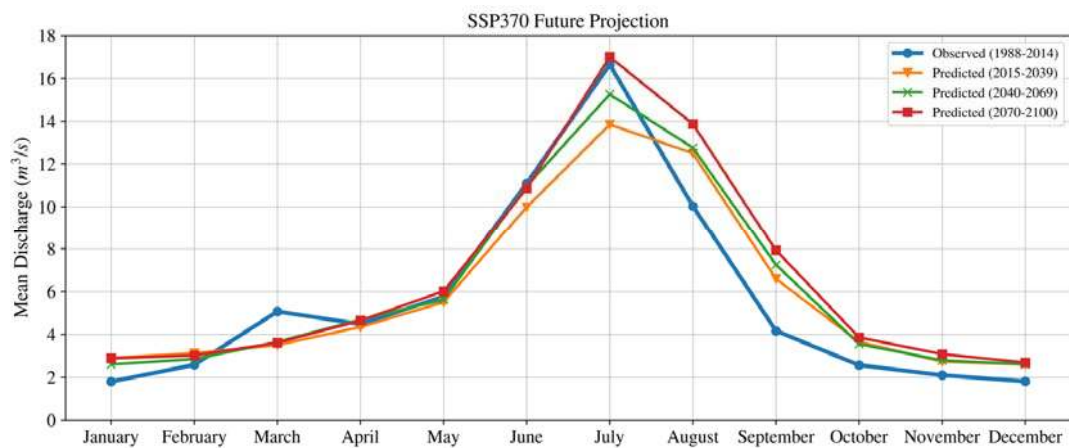


Figure 4.35 Line graph of observed and predicted mean discharge with default model parameters under SSP3-7.0 scenario

The mean discharge between 2015 and 2100 periods is depicted in SSP5-8.5, Figure 4.36. The graph indicates that the 2015–2039 period predicted a lower amount of discharge, $5.864 \text{ m}^3/\text{s}$, while the 2070–2100 period had a higher mean discharge, $6.668 \text{ m}^3/\text{s}$.

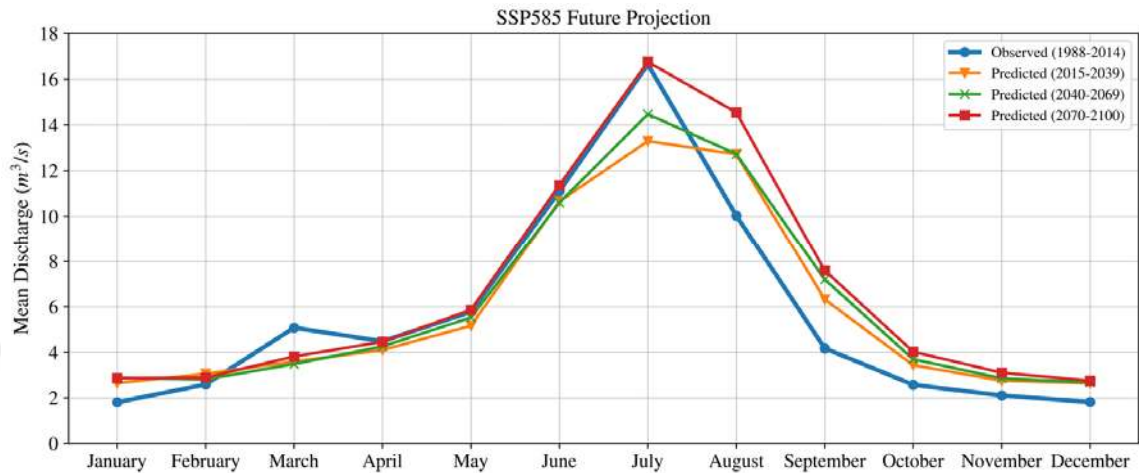


Figure 4.36 Line graph of observed and predicted mean discharge with default model parameters under SSP5-8.5 scenario

The line graph in Figure 4.37 compares the observed mean discharge with the mean discharge predicted under various SSPs from 2015 to 2039. SSP3-7.0 forecasts a higher mean discharge of $5.945 \text{ m}^3/\text{s}$ during this time frame, whereas SSP5-8.5 forecasts a lower mean discharge of $5.864 \text{ m}^3/\text{s}$.

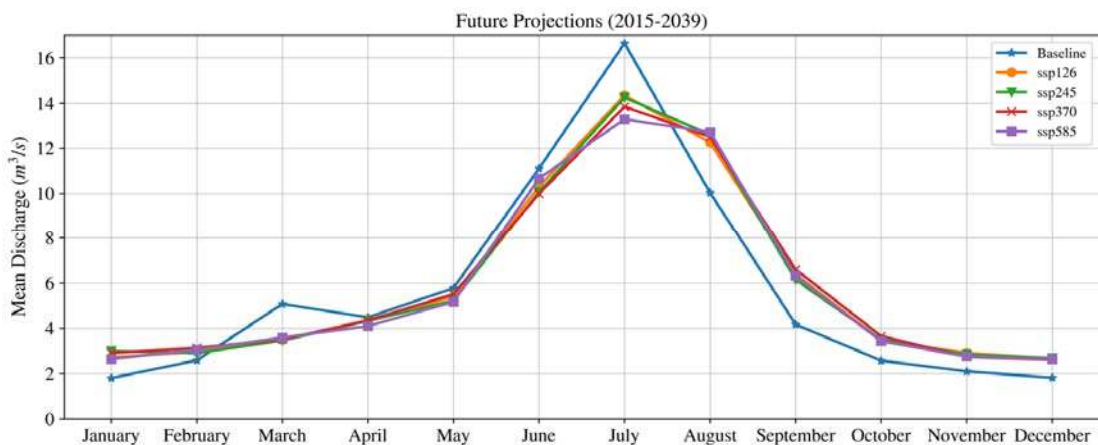


Figure 4.37 Comparison of the mean discharge between 2015-2039 years projections in CatBoost model under different SSP scenarios.

Figure 4.38 compares the observed mean discharge with the expected mean discharge under SSPs from 2040 to 2069. The scenarios exhibit similarities with one another, despite their very slight differences. With SSP3-7.0, the mean discharge amount is expected to be higher at 6.221 m³/s, whereas SSP1-2.6 has a lower amount at 6.059 m³/s.

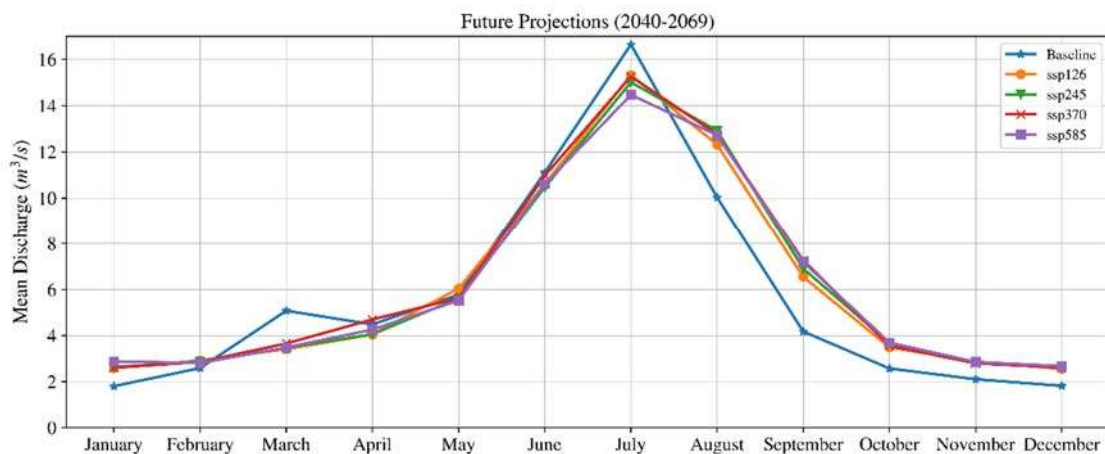


Figure 4.38 Comparison of the mean discharge between 2040-2069 years projections in CatBoost model under different SSP scenarios.

The line graph in Figure 4.39 compares the observed mean discharge with the mean discharge predicted under various SSPs between 2070 and 2100. Every SSP from January through June and October through December is parallel. The higher value of 6.668 m³/s is predicted by SSP5-8.5. and SSP1-2.6 predicts a lower value of 5.968 m³/s.

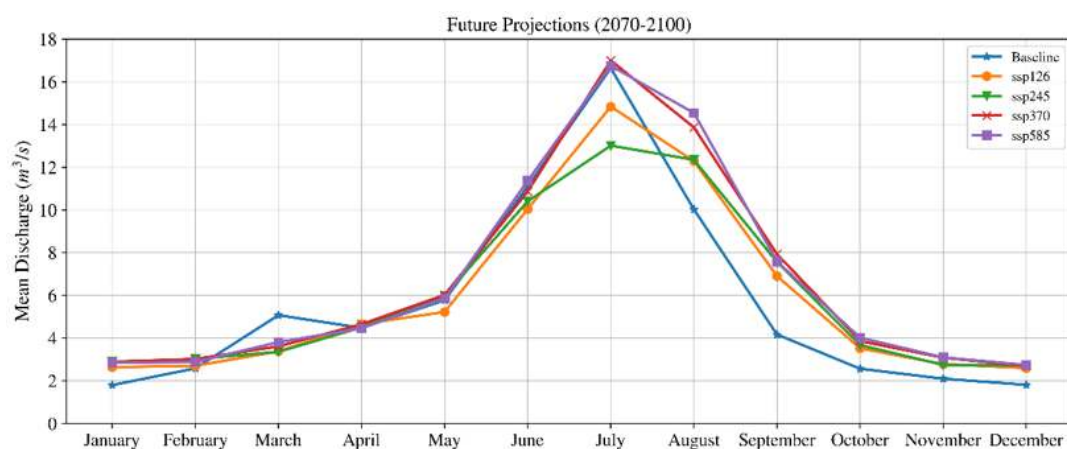


Figure 4.39 Comparison of the mean discharge between 2070-2100 years projections in CatBoost model under different SSP scenarios.

4.4 Determination of Energy Production with Predicted Discharges

It is clear from Table 4.10 that, in accordance with the observed discharge values, the average flow amounts anticipated under each scenario for the years 2015–2100 will be higher. It would be inaccurate, nevertheless, to claim that the energy derived from historical data is less than the energy anticipated for each scenario and time period in the energy calculation based on the dam heights shown in Table 4.10, which was derived using the average monthly flows. It is noted that the estimated energy amount in the SSP1-2.6 and SSP2-4.5 scenarios is lower, particularly in the 2015–2039 timeframe. The amount of energy computed under the SSP3-7.0 and SSP5-8.5 scenarios, however, shows a notable increase. The comparison of forecasting energy generation under SSPs is shown in Figure 4.40.

Table 4.10 Results of mean discharges and energy prediction

	1988-2014		2015-2039		2040-2069		2070-2100	
	Energy		Energy		Energy		Energy	
	Q(m ³ /s)	(MW)	Q(m ³ /s)	(MW)	Q(m ³ /s)	(MW)	Q(m ³ /s)	(MW)
SSP1-2.6			5.936	0.072	6.059	0.074	5.968	0.073
SSP2-4.5			5.916	0.099	6.088	0.102	6.025	0.101
SSP3-7.0			5.945	0.189	6.221	0.197	6.627	0.21
SSP5-8.5			5.864	0.232	6.097	0.241	6.668	0.263
Observed	5.681	0.12						

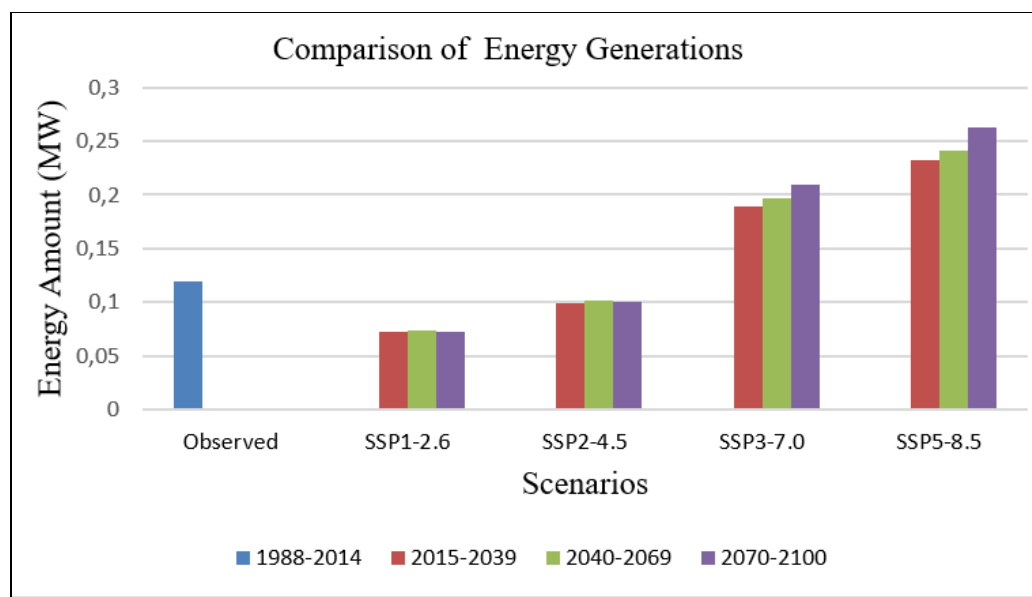


Figure 4.40 Comparison of energy generation results.

4.5 Discussion

This study shows that the CatBoost model exhibits significantly better performance than the GEP, SVM, and Ridge Regression models for discharge during the testing period (2007–2014). The future predictions of discharge data from the models have also been evaluated under four distinct scenarios: SSP1-2.6, SSP2-4.5, SSP3-7.0, and SSP5-8.5. The models made projections for the period of 2015–2100, and their performance between 2007 and 2014 was evaluated against observed values to measure the efficacy of the models. The CatBoost model, under all scenarios forecast discharges with a remarkable increase in discharge. For scenario SSP1-2.6 5.404%; SSP2-4.5 5.791%; SSP3-7.0 10.262%; and SSP5-8.5 9.312% increasing predicted. Table 4.11 shows the comparison of mean monthly discharge (m^3/s) change in future projections observed by means of (%) percentage.

Table 4.11 Comparison of mean monthly discharge (m^3/s) change in future projections observed by means of (%) percentage

	Observed	SSP1-2.6	SSP2-4.5	SSP3-7.0	SSP5-8.5
2015-2039		5.936	5.916	5.945	5.864
2040-2069		6.059	6.088	6.221	6.097
2070-2100		5.968	6.025	6.627	6.668
1988-2014	5.681				
Average	5.681	5.988	6.010	6.264	6.210
Difference (%)		5.404	5.791	10.262	9.312

The projection of annual discharge under NorESM2-MM scenarios with SSP1-2.6, SSP2-4.5, SSP3-7.0, and SSP5-8.5 by the CatBoost model on a normal scale with a trendline is given in Figure 4.42, Figure 4.43, Figure 4.45, and Figure 4.45, respectively. Out of trendline data can be considered the extreme climate condition. such as over precipitation.

It is seen from Figure 4.41 that projected discharge data has a balance in the trendline. For the periods of 2025, 2058, and 2088, the annual discharge data is out of trendline.

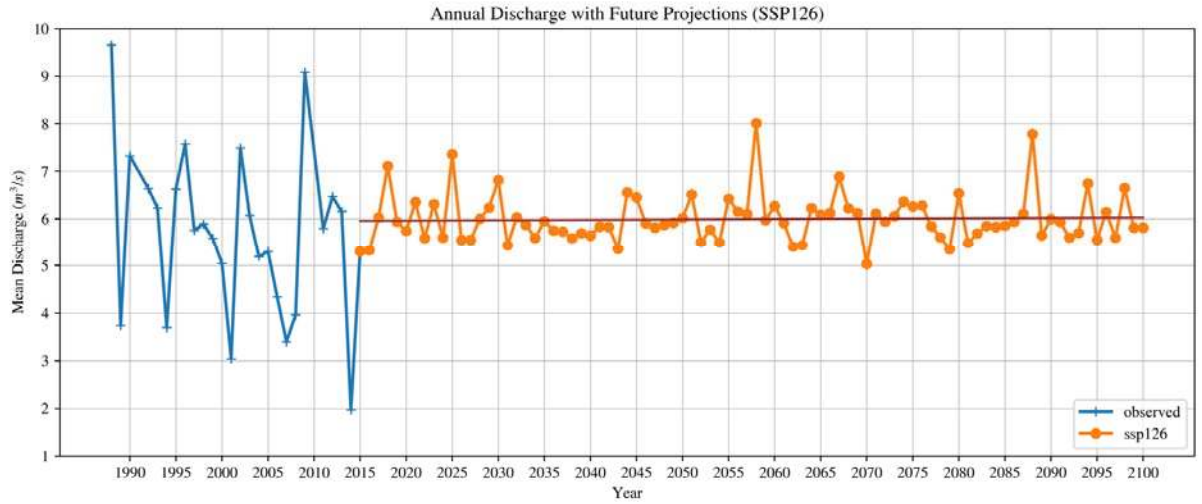


Figure 4.41 The projection of annual discharge (Q) values under the NorESM2-MM SSP1-2.6 scenario by CatBoost model and observed data between 1988-2014 in normal scale with a trendline

Figure 4.42 shows that the projected discharge is balanced with the trendline except for the periods of 2031 and 2041.

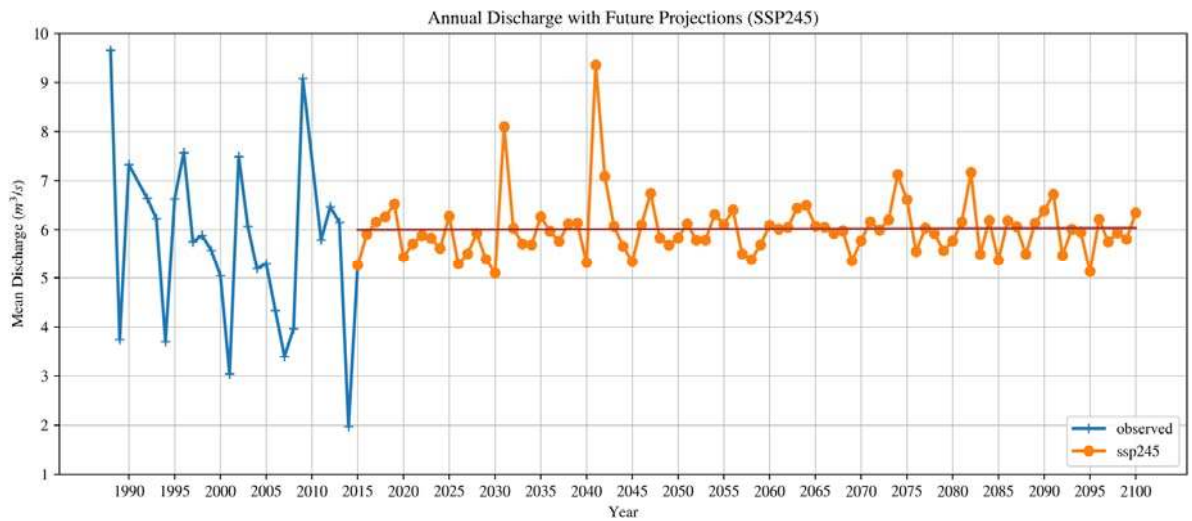


Figure 4.42 The projection of annual discharge (Q) values under the NorESM2-MM SSP2-4.5 scenario by CatBoost model and observed data between 1988-2014 in normal scale with a trendline

It is seen from the Figure 4.43 that projected discharge data has a remarkable increasing trend. For the periods of 2052, 2070, and 2096, the annual discharge data is out of trendline.

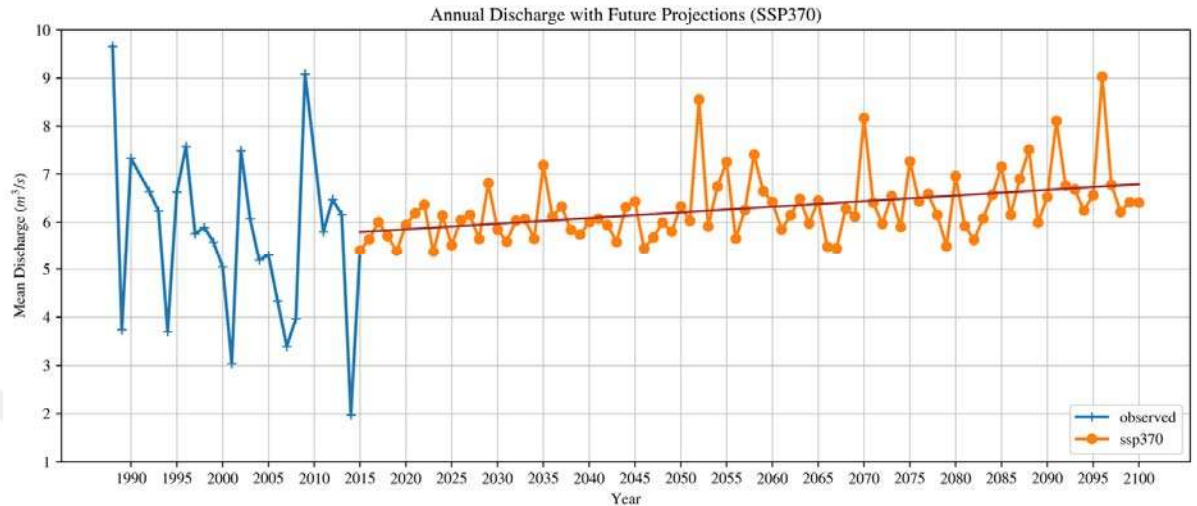


Figure 4.43 The projection of annual discharge (Q) values under the NorESM2-MM SSP3-7.0 scenario by CatBoost model and observed data between 1988-2014 in normal scale with a trendline

It is understood from Figure 4.44 that projected discharges data has a substantial increasing trend. In 2073, the annual discharge data are out of trendline.

Additionally, for periods 2015-2039 and 2040-2069, SSP3-7.0 predicts a higher monthly mean discharge of $5.945 \text{ m}^3/\text{s}$ and $6.221 \text{ m}^3/\text{s}$, respectively, while SSP5-8.5 for periods 2070-2100 predicts more discharge with $6.668 \text{ m}^3/\text{s}$. Furthermore, SSP1-2.6 predicts the less discharge with $6.059 \text{ m}^3/\text{s}$ and $5.968 \text{ m}^3/\text{s}$ for 2040-2069 and 2070-2100 periods, respectively, while SSP5-8.5 predicts $5.864 \text{ m}^3/\text{s}$ for 2015-2039.

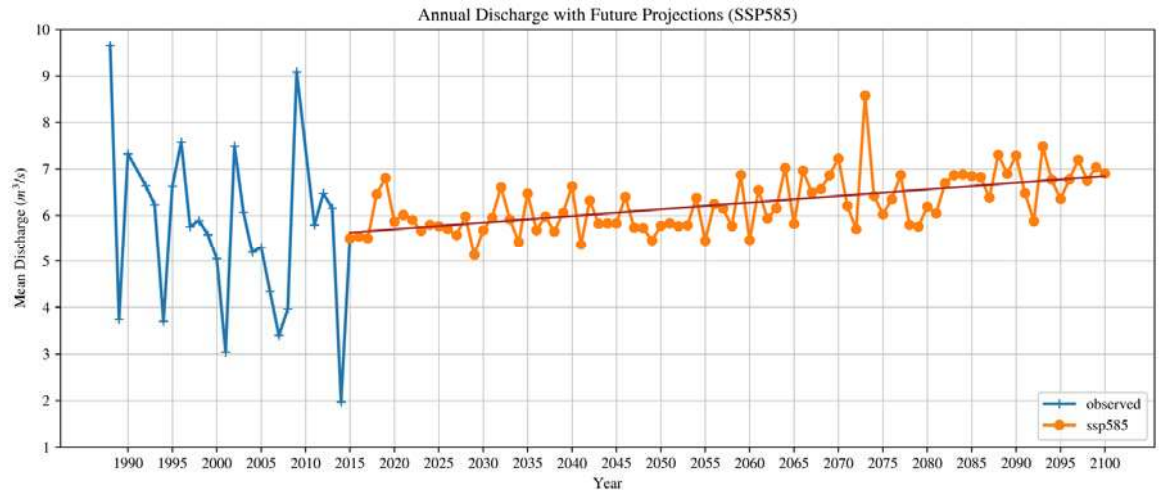


Figure 4.44 The projection of annual discharge (Q) values under the NorESM2-MM SSP5-8.5 scenario by CatBoost model and observed data between 1988-2014 in normal scale with a trendline

For each scenario, it has been calculated in this study that the predicted mean discharge is greater than the observed flow. We have determined that the current amount increases by 5.404% for SSP1-2.6 scenarios, 5.791% for SSP2-4.5, 10.262% for SSP3-7.0, and 9.312% for SSP5-8.5. It is predicted that the largest increase in the average current will occur under the SSP3-7.0 scenario with a temperature rise of 3.5–4 degrees, while the smallest increase is expected to occur under the SSP1-2.6 scenario with a temperature rise of 2 degrees. The melting of snow that will occur with a temperature increase of 3.5–4 degrees should also be considered as one of the reasons for this increase.

Several studies predicting an increase in discharge or precipitation under climate change are presented below.

In their global study, Arnell & Gosling (2013) stated that flow regimes could increase in some regions. An increase in flow is predicted, especially in high latitudes and regions where increased rainfall is expected.

Güven et al. (2024) reported in their projections for a sub-basin of the Euphrates River that increases in both flow and sediment yield are possible in the future depending on climate change scenarios. Results show;

- A 10.262% increase in the SSP3-7.0 scenario,

- A 9.312% increase in the SSP5-8.5 scenario was observed.

This result particularly indicates that changes in precipitation patterns and snowmelt, especially with increasing temperatures, can increase discharge.

In addition to rising temperatures and precipitation, the study "Predicting future impacts of climate and land use change on streamflow in the middle reaches of China's Yellow River" predicts an annual discharge increase of 49.2-115.1% over a 30-year future period, compared to the 1989-2018 period (Ma et al., 2024).

In the article titled "The Effects of Climate Change on Streamflow, Nitrogen Loads, and Crop Yields in the Gordes Dam Basin, Turkey," it is stated that the average annual flow into the reservoir is expected to increase between 0.7 and 4 m³/s during the period 2031-2060 under RCP 4.5 and 8.5 scenarios (Özdemir et al., 2024).

In the study "Impacts of climate change on streamflow of Qinglong River, China," an increase of approximately 30% in the annual average discharge is projected under the RCP4.5 and RCP8.5 scenarios (Liu & Tang, 2024).

IPCC reports (IPCC, 2021) and CMIP6 comparisons also emphasise that there will be increased river flow in some regions (particularly high latitudes in the Northern Hemisphere, basins dependent on snowmelt).

The study's most notable results include the increase in streamflow predicted under the SSP scenarios when the CatBoost model is used for simulation. Prior studies have generally suggested that increased evapotranspiration, changed precipitation patterns, and the depletion of snow reserves would result in decreased river discharge as a result of climate change. On the other hand, the current findings show the opposite trend. This phenomenon has a number of explanations.

First, the dynamics of the local climate seem to be a determining factor. According to the NorESM2-MM climate projections, precipitation over the basin is expected to increase in intensity and irregularity rather than decrease in the future, especially under high-emission pathways like SSP3-7.0 and SSP5-8.5. Because of these changes in rainfall patterns, there is more direct runoff, which raises streamflow levels.

Second, the observed increases are caused by snowmelt processes that are impacted by warming temperatures. The scenarios show a significant increase in seasonal temperatures, which speeds up snowmelt in the catchment's headwaters. This extra surface runoff, which is particularly noticeable in the spring and early summer, increases the mean discharge for the year and makes up for any losses from evapotranspiration.

Third, it's important to highlight the CatBoost algorithm's superior learning ability. CatBoost incorporates several nonlinear interactions across a broad range of meteorological predictors, in contrast to traditional models. The model reveals subtle climate signals that may have been missed by linear approaches by capturing hidden dependencies, such as the interaction between pressure anomalies, large-scale circulation indices, and precipitation patterns. Consequently, when compared to conventional hydrological models, CatBoost indicates a discharge trajectory that is more hopeful.

Fourth, the negative effects of climate change may also be lessened by basin-specific hydrological features. The hydrogeological conditions in the study area are comparatively resilient, with evapotranspiration stress being balanced by infiltration capacity, groundwater recharge, and surface retention processes. Under climate change, this hydrological buffering effect promotes a more consistent or even rising streamflow response.

Lastly, the consequences for the production of hydropower are significant. The anticipated rise in discharge under SSP3-7.0 and SSP5-8.5 results in a significant improvement in the potential for energy production. Such results highlight the significance of localised assessments, even though they run counter to the largely negative global narrative regarding climate change and water scarcity. In this instance, climate change might present chances for the basin to have better access to water and energy security.

The findings suggest that the effects of climate change on hydrology are not universally applicable across geographical boundaries. Together, increased precipitation, faster snowmelt, basin resilience, and CatBoost's sophisticated predictive capabilities account for the surprising increase in discharge. These results

demonstrate that basin-specific analyses are necessary and validate that machine learning techniques can reveal hidden patterns that alter our perception of how climate change affects water resources.



CHAPTER 5

CONCLUSIONS

5.1 Imputation of Missing Data

After imputation, Table 4.1 shows that the calculated and imputed values using CatBoost demonstrated improvements in RMSE and MAE of 15.47% and 17.82%, respectively, for the trial period, and 17.15% and 11.03%, respectively, for the test period. There has been a noticeable improvement in accuracy in the imputation carried out using CatBoost. In the SHAP graphs, it has been determined to what extent and in which direction each input affects the prediction of the missing data. Thus, the importance of the inputs that are effective in the predictions has been revealed. Before imputation, the most important role in the prediction belonged to the mean sea level pressure (mslp) input, whereas after imputation, it was understood that the most valuable input was 500 hPa Geopotential (p500).

As can be understood from the Figure 11, the air temperature at 2 m (temp) data was the second most important input, but afterwards, in Figure 12 it ranks last in importance. As it can be concluded from this study, it is effective to use CatBoost to find missing data and compare results using SHAP values in hydrological research. CatBoost's predictions become interpretable when paired with SHAP values, which enables analysts to comprehend each feature's contribution even in the case of missing data. Researchers can gain a better understanding of hydrological processes by using SHAP values to quantify the combination of SHAP values' interpretability and CatBoost's ability to handle missing data offers a useful approach to enhancing hydrological analysis, which will enhance sustainability in water resource management and enable better decision-making.

5.2 Comparison of Conventional Models and CatBoost

The goal of this study was to compare the performance of conventional models and the CatBoost algorithm, as well as to analyze discharge future projections under

climate change effects in basin area using the best model. The basin's discharge data were computed utilizing the statistical downscaling approach with GEP and SVM, CatBoost, and Ridge Regression. There has been a comparison between these methods. The investigation demonstrates that CatBoost outperforms other models with $NSE=0.636$.

The CatBoost model's projection results under the SSPs scenario were divided into three periods: 2015-2039, 2040-2069, and 2070-2100, and the monthly average discharge was graphically contrasted. The results demonstrate that SSP1-2.6 and SSP2-4.5 forecast a balancing trend, while SSP3-7.0 and SSP5-8.5 scenarios predict an increase trend between 2015 and 2100.

Considering that the monthly mean discharge of the observed data is $5.681 \text{ m}^3/\text{s}$, it indicates that in all scenarios, the predicted mean discharge for all three periods is greater than $5.681 \text{ m}^3/\text{s}$, suggesting that the measuring station and its basin are not significantly affected by climate change and that the flow is showing a positive increase.

These findings have important implications for the basin's long-term destiny. Climate change is anticipated to increase discharge in basins. These changes will affect a range of river-dependent activities, including hydropower generation. Increased discharge may improve the basin's hydropower potential and impact energy output. The proposed models' usefulness is limited by the data used in the study area. This study contributes to the field by analyzing several models under various climate change scenarios for predicting climatic data set. It is expected that the study's findings will serve as a reference for hydrologists who estimate the amount of water in river basins as well as a guidance for decision-makers in government organizations.

5.3 Comparison of Energy Production

In this study, streamflows projected to occur between 2015 and 2100 under the IPCC's future climate change scenarios were predicted using the CatBoost machine learning technique with the help of CMIP6 and NorESM2-MM variables, based on historical streamflow values measured between 1988 and 2014. When making the predictions, the years between 2015 and 2100 were divided into three periods: 2015-

2039, 2040-2069, and 2070-2100. The changes in each period for each scenario were compared with each other and with the observed flow values. According to the results obtained, the first thing to note is that the predicted flow under each period and scenario is greater than the average flow value of 5.681 m³/s measured between 1988 and 2014. This also indicates that the basin and the region where the station is located will be positively affected by climate change in terms of energy production, considering the amount of energy calculated based on the flow value.

According to each evaluated period and scenario, the highest average flow rates among the predicted flows are 5.945 m³/s for the SSP3-7.0 scenario between 2015-2039, 6.221 m³/s for the SSP3-7.0 scenario between 2040-2069, and 6.668 m³/s for the SSP5-8.5 scenario between 2070-2100. The estimated smallest average flow rates are associated with the SSP2-4.5 scenario at 5.916 m³/s between 2015-2039, the SSP1-2.6 scenario at 6.059 m³/s between 2040-2069, and the SSP1-2.6 scenario at 5.968 m³/s between 2070-2100.

While calculating the impact of the obtained flows on energy production, many complex data such as evaporation, HDD and CDD values, spillway quantity, dam volume, etc., were not taken into account. Although it is possible to perform more complex calculations for a more realistic and reliable result, in this study, only flow-related changes were predicted at a simpler level.

The estimated energy amounts in MW based on the obtained average flow data are presented in Table 4.10, and as can be seen from the table, when compared to the 0.12 MW of energy that can be produced based on the observed flow, it is evident that there will be a significant increase in energy amounts between 2015-2100, especially under the SSP3-7.0 and SSP5-8.5 scenarios. Table 4 shows the percentage differences between the amount of energy that can be produced under climate change from 2015 to 2100 and the amount of energy that can be obtained based on the observed flow value from 1988 to 2014. The largest increase will be 119.167 % under the SSP5-8.5 scenario in the 2070-2100 period.

Table 5.1 Comparison of energy production differences (%) under SSP scenarios based on observed discharges

	2015-2039	2040-2069	2070-2100
	%	%	%
SSP1-2.6	-40.000	-38.333	-39.167
SSP2-4.5	-17.500	-15.000	-15.833
SSP3-7.0	57.500	64.167	75.000
SSP5-8.5	93.333	100.833	119.167

Increasing a river's or water source's discharge generally raises the amount of water that can power turbines, which in turn increases the potential for hydropower generation. As more water passes through the system, more kinetic energy is available to be transformed into electrical energy, which can lead to elevated energy production. However, environmental factors like ecosystem health and sediment transport, as well as reservoir capacity and turbine efficiency, also have an impact on hydropower potential overall.

The data used in the study area limits the usefulness of the suggested models. By examining multiple models under diverse climate change scenarios for climatic data set prediction, this study advances the field. The study's conclusions are anticipated to be used as a guide for government decision-makers and as a reference by hydrologists who calculate the volume of water in river basins.

REFERENCES

- Abdulhussain, S. H., Mahmmud, B. M., Naser, M. A., Alsabah, M. Q., Ali, R., & Al-Haddad, S. A. R. (2021). A robust handwritten numeral recognition using hybrid orthogonal polynomials and moments. *Sensors*, 21(6), 1999.
- Arnell, N. W., & Gosling, S. N. (2013). The impacts of climate change on river flow regimes at the global scale. *Journal of Hydrology*, 486, 351–364. <https://doi.org/10.1016/j.jhydrol.2013.02.010>
- Bentsen, M., Bethke, I., Debernard, J. B., Iversen, T., Kirkevåg, A., Seland, Ø., ... & Kristjánsson, J. E. (2013). *The Norwegian Earth System Model, NorESM1-M–Part 1: Description and basic evaluation of the physical climate*, *Geosci. Model Dev.*, 6, 687–720.
- Bonakdari, H., Ebtehaj, I., Gharabaghi, B., Sharifi, A., & Mosavi, A. (2021). Prediction of discharge capacity of labyrinth weir with gene expression programming. In *Intelligent Systems and Applications: Proceedings of the 2020 Intelligent Systems Conference (IntelliSys) Volume 1* (pp. 202-217). Springer International Publishing.
- Chen, B., Chen, Z., Song, C., & Song, Y. (2024). Integrated forecasting method of medium-and long-term runoff by ridge regression based on optimal sub-model selection. *Water Supply*, 24(3), 799-811.
- Cortes, C., & Vapnik, V. (1995). Support-vector networks. *Machine learning*, 20(3), 273-297.
- Coskun, S., & Akbas, A. (2024). Revealing the future complexity of urban water scarcity and drought via support vector machine: Case from semi-arid Bursa urban area. *Urban Climate*, 58, 102211.
- Dalla Torre, D., Lombardi, A., Menapace, A., Zanfei, A., & Righetti, M. (2024). Exploring the feasibility of Support Vector Machine for short-term hydrological forecasting in South Tyrol: challenges and prospects. *Discover Applied Sciences*, 6(4), 154.
- Eker, R., Alkiş, K. C., Uçar, Z., & Aydın, A. (2023). Ormancılıkta makine öğrenmesi kullanımı. *Turkish Journal of Forestry*, 24(2), 150-177.

- Esha, R., & Imteaz, M. A. (2020). Pioneer use of gene expression programming for predicting seasonal streamflow in Australia using large scale climate drivers. *Ecohydrology*, 13(8), e2242.
- Essam, Y., Huang, Y. F., Ng, J. L., Birima, A. H., Ahmed, A. N., & El-Shafie, A. (2022). Predicting streamflow in Peninsular Malaysia using support vector machine and deep learning algorithms. *Scientific Reports*, 12(1), 3883.
- Eyring, V., Bony, S., Meehl, G. A., Senior, C. A., Stevens, B., Stouffer, R. J., & Taylor, K. E. (2016). Overview of the Coupled Model Intercomparison Project Phase 6 (CMIP6) experimental design and organization. *Geoscientific Model Development*, 9(5), 1937–1958. <https://doi.org/10.5194/gmd-9-1937-2016>
- Ferreira, C. (2001). Gene expression programming: a new adaptive algorithm for solving problems. *Complex Systems*, 13(2), 87-129. *arXiv preprint cs/0102027*.
- Fuladipanah, M., Azamathulla, H. M., Tota-Maharaj, K., Mandala, V. & Chadee, A. (2023) Precise forecasting of scour depth downstream of flip bucket spillway through data-driven models. *Results Eng.* 20, 101604. <https://doi.org/10.1016/j.rineng.2023.101604>.
- Gözütok, S., & Özdemir, O. N. (2004). Genetik algoritma yöntemi ile su şebekelerinde hidrolik kalibrasyonun geliştirilmesi. *Gazi Üniversitesi Mühendislik Mimarlık Fakültesi Dergisi*, 19(2).
- Gupta, H. V., Kling, H., Yilmaz, K. K., & Martinez, G. F. (2009). Decomposition of the mean squared error and NSE performance criteria: Implications for improving hydrological modelling. *Journal of hydrology*, 377(1-2), 80-91.
- Güven, A., Gün, M. V., & Pala, A. (2024). Modeling and future projection of streamflow and sediment yield in a sub-basin of Euphrates River under the effect of climate change. *Journal of Water and Climate Change*, 15(6), 2628-2647.
- Hyndman, R. J., & Athanasopoulos, G. (2018). *Forecasting: principles and practice*. OTexts.
- Hoerl, A. E., & Kennard, R. W. (1970). Ridge regression: Biased estimation for nonorthogonal problems. *Technometrics*, 12(1), 55-67.
- Huang, G., Wu, L., Ma, X., Zhang, W., Fan, J., Yu, X., ... & Zhou, H. (2019). Evaluation of CatBoost method for prediction of reference evapotranspiration

- in humid regions. *Journal of Hydrology*, 574, 1029-1041.
- Hurrell, J. W., Holland, M. M., Gent, P. R., Ghan, S., Kay, J. E., Kushner, P. J., ... & Marshall, S. (2013). The community earth system model: a framework for collaborative research. *Bulletin of the American Meteorological Society*, 94(9), 1339-1360.
- Hussein, E., Ghaziasgar, M., & Thron, C. (2020, July). Regional rainfall prediction using support vector machine classification of large-scale precipitation maps. In *2020 IEEE 23rd International Conference on Information Fusion (FUSION)* (pp. 1-8). IEEE.
- IPCC (2000). Special Report on Emissions Scenarios, 2000.
<https://www.ipcc.ch/site/assets/uploads/2018/03/sres-en.pdf>
- IPCC (2014). *Climate Change 2014: Impacts, Adaptation, and Vulnerability. Part A: Global and Sectoral Aspects*. Cambridge University Press.
- IPCC (2019). Summary for Policymaker, 2019
https://www.ipcc.ch/site/assets/uploads/sites/4/2019/12/02_Summary-for-Policymakers_SPM.pdf
- IPCC (2021). *Climate Change 2021: The Physical Science Basis*. Contribution of Working Group I to the Sixth Assessment Report of the Intergovernmental Panel on Climate Change. Cambridge University Press.
<https://doi.org/10.1017/9781009157896>
- Jisha, G. (2024, August). Enhanced Weather Prediction with Feature Engineered, Time Series Cross Validated Ridge Regression Model. In *2024 Control Instrumentation System Conference (CISCON)* (pp. 1-6). IEEE.
- Karimi, S., Shiri, J., Kisi, O., & Makarynsky, O. (2012). Forecasting water level fluctuations of Urmieh Lake using gene expression programming and adaptive neuro-fuzzy inference system. *The International Journal of Ocean and Climate Systems*, 3(2), 109-125.
- Kendall, M. G. (1975). "Rank Correlation Methods." *Griffin*, 4th edition.
- Kilinc, H. C., Ahmadianfar, I., Demir, V., Heddami, S., Al-Areeq, A. M., Abba, S. I., ... & Yaseen, Z. M. (2023). Daily scale river flow forecasting using hybrid gradient boosting model with genetic algorithm optimization. *Water resources management*, 37(9), 3699-3714.

- Krause, P., Boyle, D. P., & Base, F. (2005). Comparison of different efficiency criteria for hydrological model assessment. *Advances in Water Resources*, 28(1), 89-103.
- Kundra, K. S. R., Lakshmi, B. J., Venugopal, I. V. S., & Guthula, V. (2023). Flood Prediction using MLP, CatBoost and Extra-Tree Classifier. *International Journal on Recent and Innovation Trends in Computing and Communication*, 11(7), 35-44.
- Langhammer, J. (2023). Flood simulations using a sensor network and support vector machine model. *Water*, 15(11), 2004.
- Liu, X., & Tang, Z. (2024). Impacts of climate change on streamflow of Qinglong River, China. *Journal of Water and Climate Change*, 15(1), 233-270.
- Lundberg, S. M., & Lee, S. I. (2017). A unified approach to interpreting model predictions. *Advances in neural information processing systems*, 30.
- Ma, X., Li, Z., Ren, Z., Shen, Z., Xu, G., & Xie, M. (2024). Predicting future impacts of climate and land use change on streamflow in the middle reaches of China's Yellow River. *Journal of Environmental Management*, 370, 123000.
- Milly, P. C., Dunne, K. A., & Vecchia, A. V. (2005). Global pattern of trends in streamflow and water availability in a changing climate. *Nature*, 438(7066), 347-350.
- Mishra, B. R., Vogeti, R. K., Jauhari, R., Raju, K. S., & Kumar, D. N. (2024). Boosting algorithms for projecting streamflow in the Lower Godavari Basin for different climate change scenarios. *Water Science & Technology*, 89(3), 613-634.
- Montgomery, J. M., Hollenbach, F. M., & Ward, M. D. (2012). Improving predictions using ensemble Bayesian model averaging. *Political Analysis*, 20(3), 271-291.
- Moriasi, D.N.; Arnold, J.G.; Van Liew, M.W.; Bingner, R.L.; Harmel, R.D.; Veith, T.L. Model evaluation guidelines for systematic quantification of accuracy in watershed simulations. *Trans. ASABE* 2007, 50, 885–900.
- Msiza, I. S., Nelwamondo, F. V., & Marwala, T. (2007, October). Artificial neural networks and support vector machines for water demand time series forecasting. In *2007 IEEE International Conference on Systems, Man and Cybernetics* (pp. 638-643). IEEE.

- Nash, J. E., & Sutcliffe, J. V. (1970). River flow forecasting through conceptual models part I—A discussion of principles. *Journal of hydrology*, 10(3), 282-290.
- Ogundolie, O. I., Olabiyisi, S. O., Ganiyu, R. A., Jeremiah, Y. S., & Ogundolie, F. A. (2024). Evaluation of the Performance of Categorical Boosting Algorithm for Flood Prediction in Osun River Basin. *Journal of Computing and Communication*, 3(2), 23-30.
- Önen, F., & Oral, Ş. O. (2019). Genetik İfadeli Programlama İle Taşkın Öteleme Modellemesi. *Dicle Üniversitesi Mühendislik Fakültesi Mühendislik Dergisi*, 10(1), 263-273.
- Özdemir, A., Volk, M., Strauch, M., & Witing, F. (2024). The Effects of Climate Change on Streamflow, Nitrogen Loads, and Crop Yields in the Gordes Dam Basin, Turkey. *Water*, 16(10), 1371.
- Prokhorenkova, L., Gusev, G., Vorobev, A., Dorogush, A. V., & Gulin, A. (2018). CatBoost: unbiased boosting with categorical features. *Advances in neural information processing systems*, 31.
- Radhakrishnan, R., & CA, L. D. (2023, July). Groundwater Level Prediction Using Support Vector Machine and M5 Model Tree-A Case Study. In *Proceedings of the 7th Biennial International Conference On Emerging Trends in Engineering, Science & Technology (ICETEST 2023)*.
- Raza, A., Vishwakarma, D. K., Acharki, S., Al-Ansari, N., Alshehri, F., & Elbeltagi, A. (2024). Use of gene expression programming to predict reference evapotranspiration in different climatic conditions. *Applied water science*, 14(7), 152.
- Riahi, K., Van Vuuren, D. P., Kriegler, E., Edmonds, J., O'Neill, B. C., Fujimori, S., ... & Tavoni, M. (2017). The Shared Socioeconomic Pathways and their energy, land use, and greenhouse gas emissions implications: An overview. *Global environmental change*, 42, 153-168.
- Ribeiro, M. T., Singh, S., & Guestrin, C. (2016, August). " Why should i trust you?" Explaining the predictions of any classifier. In *Proceedings of the 22nd ACM SIGKDD international conference on knowledge discovery and data mining* (pp. 1135-1144).

- Samadi-Koucheksaraee, A., & Chu, X. (2024). Development of a novel modeling framework based on weighted kernel extreme learning machine and ridge regression for streamflow forecasting. *Scientific Reports*, 14(1), 30910.
- Sharma, B., & Goel, N. K. (2024). Streamflow prediction using support vector regression machine learning model for Tehri Dam. *Applied Water Science*, 14(5), 99.
- Singh, P., Shamseldin, A. Y., Melville, B. W., & Wotherspoon, L. (2023). Development of statistical downscaling model based on Volterra series realization, principal components and ridge regression. *Modeling Earth Systems and Environment*, 9(3), 3361-3380.
- Szczepanek, R. (2022). Daily streamflow forecasting in mountainous catchment using XGBoost, LightGBM and CatBoost. *Hydrology*, 9(12), 226.
- Tikhonov, A. N. (1943). On the stability of the solution of the ill-posed problem. *Doklady Akademii Nauk*, 39(5), 195-198.
- Vapnik, V. (2013). *The nature of statistical learning theory*. Springer science & business media.
- Wang, D., & Qian, H. (2023). CatBoost-based automatic classification study of river network. *ISPRS International Journal of Geo-Information*, 12(10), 416.
- Willmott, C. J., & Matsuura, K. (2005). Advantages of the mean absolute error (MAE) over the root mean square error (RMSE) in assessing average model performance. *Climate research*, 30(1), 79-82.
- XII. National Hydrology Congress (16-19 October 2024) (Samsun Büyükşehir Belediyesi Şehit Ömer Halis Demir Çok Amaçlı Salonu) hidrolojikongresi.samsun.edu.tr
- Yaman, Y., & Önen, F. (2023). Yağış-Akış İlişkisinin GEP ve ANFIS İle Modellenmesi. *Dicle Üniversitesi Mühendislik Fakültesi Mühendislik Dergisi*, 14(3), 489-498.
- Yi, S., Yoon, J., Lee, C., Lee, S., Ji, J., Lee, E., & Yi, J. (2025). Advancing flow duration curve prediction in ungauged basins using machine learning and deep learning. *Hydrology and Earth System Sciences Discussions*, 2025, 1-31.
- Zandi, O., Nasser, M., & Zahraie, B. (2023). A locally weighted linear ridge regression framework for spatial interpolation of monthly precipitation over an

orographically complex area. *International Journal of Climatology*, 43(6), 2601-2622.



CURRICULUM VITAE

I attended primary school, middle school, and high school in Ordu. In 2012, I graduated as an honor student and top student of the department from the Environmental Engineering of Atatürk University. Between 2014 and 2017, I pursued a master's degree with a thesis in the hydraulic sub-discipline of the Civil Engineering Department at Gaziantep University and graduated with high honors. At the same time, I graduated with high honors from the Department of Occupational Health and Safety at Gaziantep University in 2023. Between 2018 and 2025, I completed my doctoral education in the same department with high honors. My work experience is as follows: Between 2014 and 2020, I worked as an environmental engineer at Turkish Electricity Transmission Corporation (TEİAŞ), focusing on hazardous and non-hazardous waste management and environmental impact assessment (EIA) reports for high voltage lines. From 2017 to 2020, I was involved in the implementation stages of Quality Management Systems at TEİAŞ. Since 2020, I have been working at the same workplace as the Chief Engineer of the Integrated Management System, and I serve as the Lead Auditor in the fields of ISO 9001, ISO 14001, ISO 45001, ISO/IEC 27001, ISO 50001, ISO 55001, and ISO 10002 management systems. At the same time, I hold an Energy Management certificate and a Class B Occupational Safety Specialist certificate, both obtained after completing the training and exam. I am married, a mother of one beautiful girl, and proficient in English.