

A machine learning bias correction method for precipitation corresponding to weather conditions using simple input data

Takao Yoshikane¹ and Kei Yoshimura²

¹The University of Tokyo

²University of Tokyo

November 23, 2022

Abstract

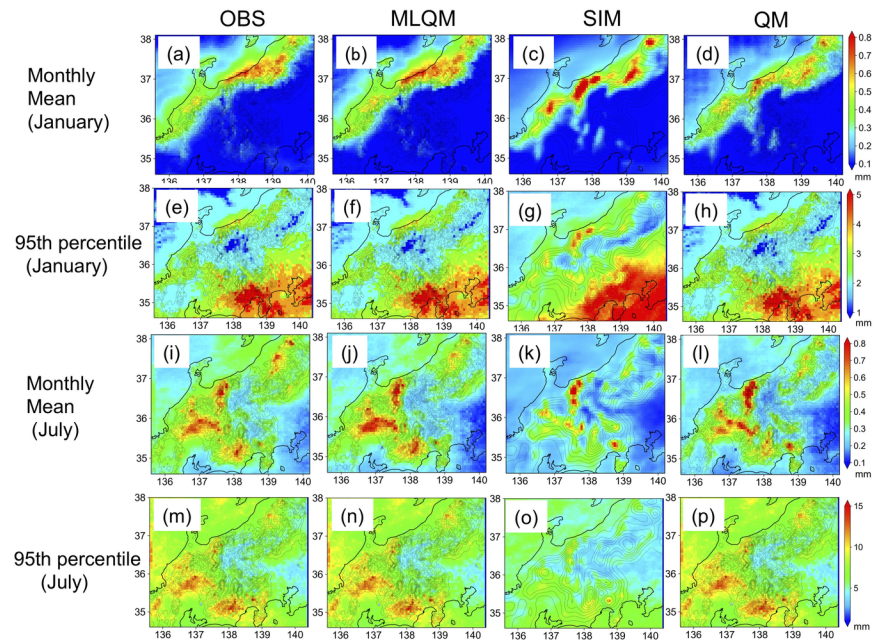
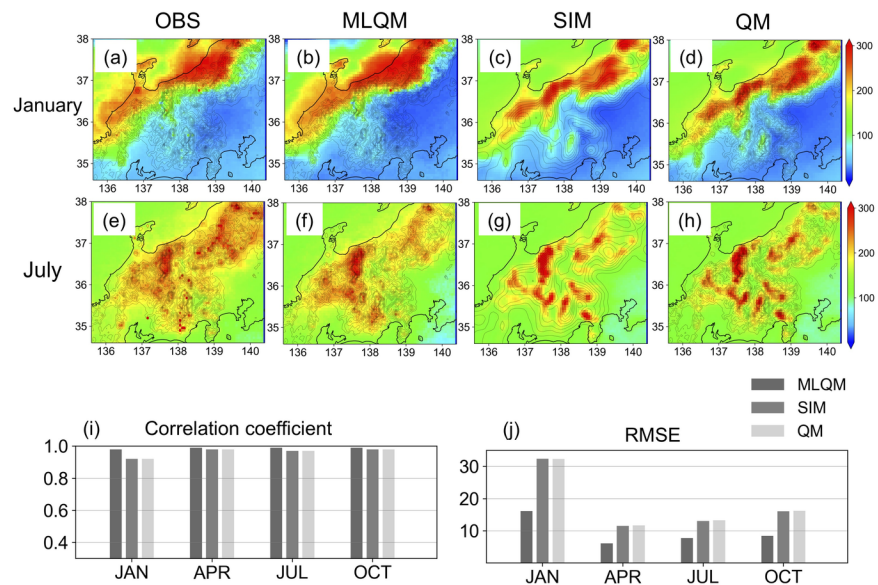
Various bias correction methods have recently been proposed using machine learning techniques. Generally, machine learning methods are fairly complicated, and it is extremely difficult to explain how machine learning corrects model biases. Accordingly, researchers perpetually seek to apply machine learning methods to diverse cases and to determine whether these methods are reliable. Here, we developed a machine learning method using simple input data by assuming a relation between observed and simulated precipitation corresponding to weather conditions. This simple method can find the optimal relation without employing dimension reduction and can facilitate the comprehension of precipitation characteristics. According to a validation experiment, this simple method can correct the precipitation frequency corresponding to the orography and estimate the local precipitation distribution characteristics, resulting in values similar to the observed data even when data are forecasted more than 24 hours from the initial time.

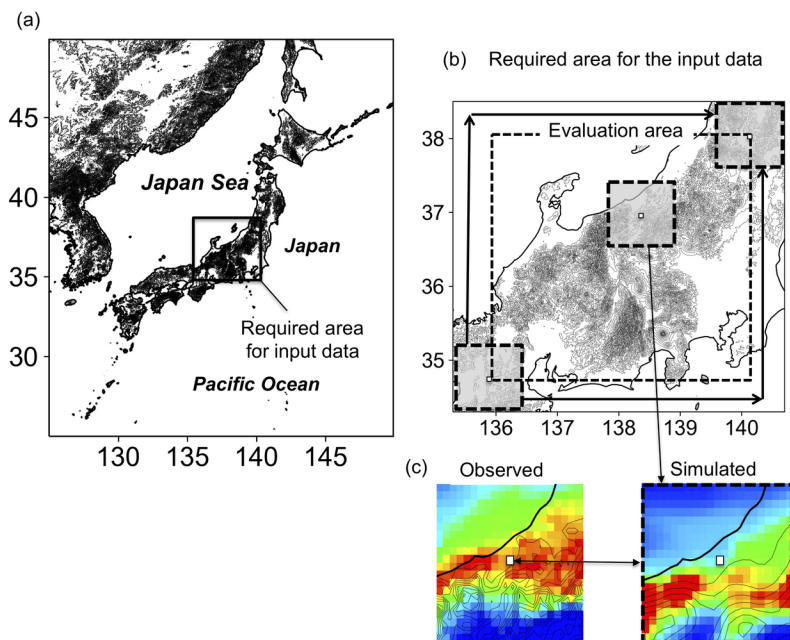
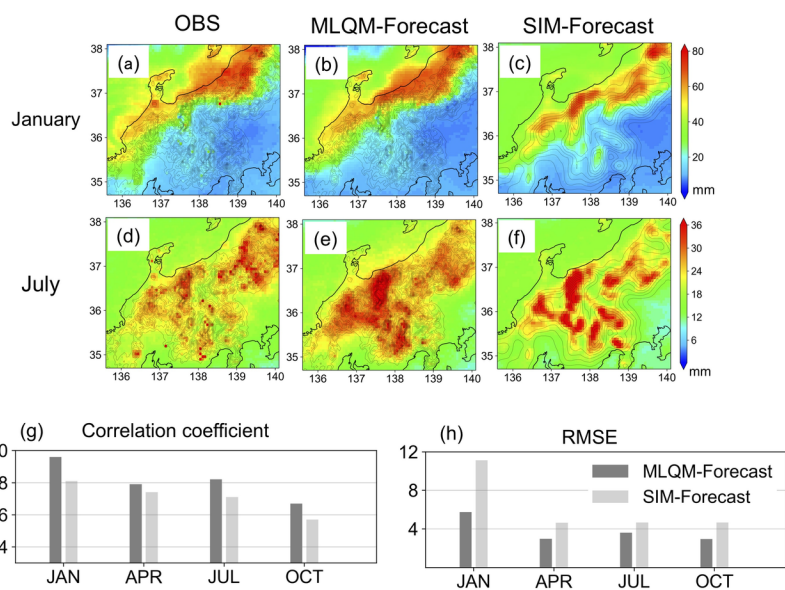
Hosted file

essoar.10507695.1.docx available at <https://authorea.com/users/545460/articles/602042-a-machine-learning-bias-correction-method-for-precipitation-corresponding-to-weather-conditions-using-simple-input-data>

Hosted file

agusupporting-information_yoshikane-20210706.docx available at <https://authorea.com/users/545460/articles/602042-a-machine-learning-bias-correction-method-for-precipitation-corresponding-to-weather-conditions-using-simple-input-data>





T. Yoshikane¹ and K. Yoshimura²

¹Institute of Industrial Science, The University of Tokyo, 5-1-5, Kashiwanoha, Kashiwa-shi, Chiba, 277-8574, Japan

²Earth Observation Research Center, Japan Aerospace Exploration Agency, 2-1-1, Sengen, Tsukuba, Ibaraki, 305-8505, Japan.

Corresponding author: Takao Yoshikane (takao-y@iis.u-tokyo.ac.jp)

Key Points:

- To developed a simple bias correction method for precipitation by assuming the relation between observed and simulated precipitation.
- The method corrects the precipitation frequency corresponding to the orography and estimate the precipitation distribution characteristics.
- Forecast data could also be corrected by recognizing the characteristics of precipitation system in an area spanning approximately 100 km².

Abstract

Various bias correction methods have recently been proposed using machine learning techniques. Generally, machine learning methods are fairly complicated, and it is extremely difficult to explain how machine learning corrects model biases. Accordingly, researchers perpetually seek to apply machine learning methods to diverse cases and to determine whether these methods are reliable. Here, we developed a machine learning method using simple input data by assuming a relation between observed and simulated precipitation corresponding to weather conditions. This simple method can find the optimal relation without employing dimension reduction and can facilitate the comprehension of precipitation characteristics. According to a validation experiment, this simple method can correct the precipitation frequency corresponding to the orography and estimate the local precipitation distribution characteristics, resulting in values similar to the observed data even when data are forecasted more than 24 hours from the initial time.

Plain Language Summary

Supervised machine learning methods require appropriate training data. If the training data are inappropriate, machine learning cannot correctly estimate the distribution. In general, appropriate objective (observed) and explanatory (simulated) training data values are necessary to reduce the model bias. However, it is difficult to find the relation between observed and simulated data because we cannot determine how well a numerical model can represent real phenomena; the method may be useful in a specific case but inapplicable otherwise. Moreover, if we do not understand the relation sufficiently, the reliability decreases considerably. Therefore, a simple, interpretable relation is required to resolve the issue and improve the method. Here, we developed a bias correction method using machine learning by assuming a simple relation between observed

and simulated precipitation. We confirmed that this method can modify the precipitation frequency to produce values similar to the observed precipitation data. By applying a hypothesis-verification approach, we expect to estimate the behavior of the machine learning method more simply and improve its reliability.

1 Introduction

Various bias correction methods using machine learning (ML) technique have been rapidly developed in recent years. However, ML approaches are generally quite complicated, and their outputs can be extremely difficult to understand (Castelvecchi, 2016; Rudin, 2019). In ML, both objective and explanatory variables (feature vectors) are necessary to correct the bias. However, it is quite difficult to determine the relation between the observations and simulated outputs (Michelangeli et al., 2009). Therefore, to obtain optimal inputs before applying ML methods, some studies have utilized dimension-reduction techniques, of which there are several kinds, such as filter and wrapper methods that use statistical analyses (Hall et al., 1999). In addition, deep learning has recently been used to automatically select optimal feature vectors (LeCun et al., 2015).

In any case, because the relations between observations and simulated outputs are extremely complicated, explaining why certain feature vectors are selected remains quite difficult. The selected feature vectors that adapt well to one case might not be generalizable to other cases. Consequently, solving the above-mentioned problems and improving bias correction methods are considerably difficult, and simple ML methods without complicated dimension reductions are required to reasonably understand the relations between observations and simulated outputs.

Local precipitation is caused by the complicated interactions among large-scale atmospheric fields; local factors, such as complex orography; and mesoscale convective systems. Generally, it is difficult to accurately simulate local precipitation due to the incompleteness of numerical models and parameterizations. Therefore, the frequencies of simulated precipitation are largely different from those of observed precipitation. Consequently, long-term simulated precipitation distributions contain large errors, and bias correction is required to correct the simulated frequencies and amounts of local precipitation.

Nevertheless, the mesoscale model of the Japan Meteorological Agency (JMA) can effectively represent the temporal variations in precipitation in a large area caused by cold and warm fronts (Saito et al., 2006). Therefore, some relations between the observed and simulated precipitation are assumed through large-scale weather conditions. ML can estimate the characteristics of precipitation by reducing the model bias if the linkage between the observed and simulated precipitation can be recognized. Here, we estimate local precipitation by identifying linkages with ML based on an assumed linkage between the distributions of the observed precipitation and the simulated precipitation. On the other hand, the characteristics of precipitation simulated far from the initial time might be

considerably different from the characteristics of precipitation simulated immediately after the data assimilation. Therefore, we verify the effectiveness of this method to improve the precipitation-estimation performance as well as its applicability for forecasting precipitation data.

2 Methods

2.1 Bias correction of precipitation using machine learning

The target area for the method validation is shown in Figure 1. The water resources in this area are necessary for numerous extensive urban areas in Japan, such as Tokyo. Moreover, this area has suffered from many water disasters. Therefore, it is important to accurately estimate the precipitation distribution in this region. The observed and simulated precipitation values were obtained from the Radar-Automated Meteorological Data Acquisition System (AMeDAS) (Makihara et al., 1996) and the mesoscale-model grid point values (MSM-GPV) dataset (Ishikawa and Koizumi 2002; JMA, 2019; Saito et al. 2006), respectively. The observed data were modified to a resolution of 0.06 degrees to correspond to the grid size of the simulated precipitation. The simulated precipitation distribution characteristics are substantially different from the observed precipitation distribution characteristics due to topographic differences. Moreover, as shown in Figure S1, the MSM-GPV data can represent the characteristics of the temporal variations in the area-averaged precipitation intensity over a wide area corresponding to large-scale weather patterns. Therefore, the distribution patterns of the simulated precipitation over a wide area are assumed to be connected with the observed precipitation distribution through weather patterns.

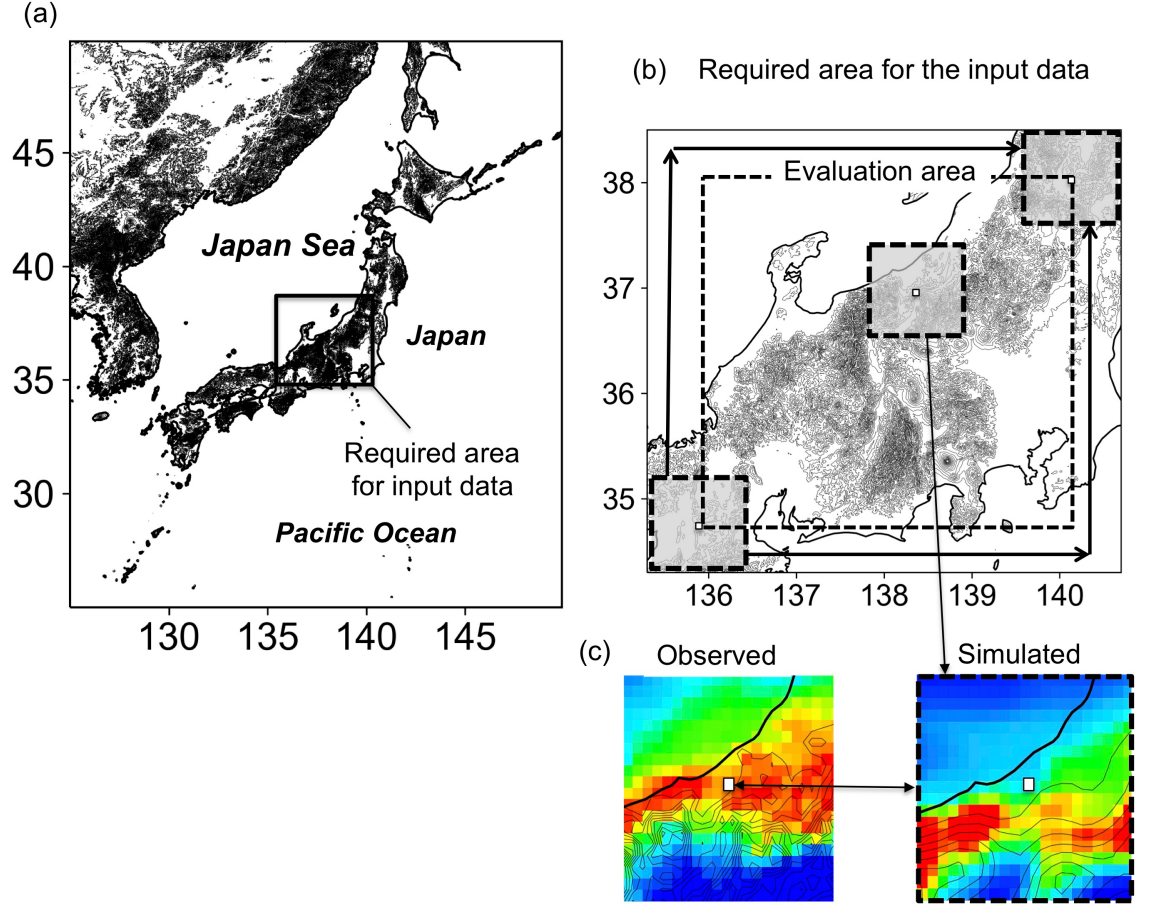


Figure 1. The target domain of this study. a. The area surrounding the target domain. b. The required area for the input data and the evaluation area of the target domain. c. Twelve-year average monthly mean observed and simulated precipitation in January; the simulated precipitation (explanatory variables) obtained in the required area were employed as the input data.

Based on the above assumption, we used the simulated precipitation over the 21×21 grid cells (almost $113 \text{ km} \times 113 \text{ km}$) covering the area shown in Figure 1 as the explanatory variable (feature vector). The observed precipitation, which was measured at the center of the area of the explanatory variable, was used as the objective variable. Therefore, the area required to establish the explanatory variable was larger than the evaluation area (Figure 1b). The ML method produces a classifier using a pair of simulated precipitation (explanatory variable) and observed precipitation (objective variable) at each grid point. In the training data, the observed precipitation corresponds to the simulated precipitation from the initial time of data assimilation, which was conducted every 3 hours,

to one hour following the assimilation. We used the observed and simulated precipitation values over 11 years (from 2007 to 2018, excluding the target year) as the training dataset and evaluated the estimated local precipitation in the area. We confirmed the generalizability of the method by conducting a cross-validation using the estimated precipitation from 2007 to 2018. An overview of the method used to estimate local precipitation is shown in Figure S2.

We used a support vector machine regression model (SVM-SVR) (Smola & Schölkopf, 2004) as the method in this study. A SVM is a supervised learning method based on part of a dataset in which predictions are obtained with a support vector. A SVM attempts to obtain the optimal results by finding the maximum-margin hyperplane, which is determined by maximizing the distance between the support vectors. Previous studies have indicated the advantages and disadvantages of SVMs relative to other ML methods, such as neural networks and random forests (Al-Anazi, & Gates 2012; Cherkassky et al., 2004; Liu et al., 2017; Sivapragasam et al., 2001). For example, SVR has been shown to perform well with small sample sizes (Al-Anazi, & Gates 2012). Thus, this method would also be useful for recognizing rare precipitation events with small sample sizes. SVMs have been employed in various fields, such as meteorology, hydrology, disasters, and water resources (Chen et al., 2019; Fan et al. 2018; Sachindra et al., 2018). The support vector machine library in the scikit-learn system (Epsilon-Support Vector Regression in scikit-learn 0.24.2) (Pedregosa et al., 2011) was used in this study.

2.2 Determination of hyperparameters

The SVR method requires the gamma, C, and epsilon hyperparameters to be configured. Gamma is a kernel function parameter that specifies the width of the Gaussian radial basis function (RBF) kernel, whereas C is the penalizing constraint error and epsilon is the width of the insensitive zone (Smets et al., 2007). Determining these hyperparameters is very important for improving the generalizability of the precipitation estimations. However, substantial computational resources are necessary to determine the optimal parameters (Anguita et al., 2010; Cherkassky & Ma, 2004). Therefore, it is necessary to obtain the optimal hyperparameters effectively. The hyperparameters could be configured at each point in the method; however, this approach is extremely inefficient because considerable computational resources are required to determine the optimal values in the entire domain. Therefore, we applied the specified hyperparameter values to all grid cells in the domain according to the following procedure. We first estimated the optimal hyperparameter values based on a random search (Bergstra & Bengio, 2012) on some grid points in the domain. The optimal values of gamma, C, and epsilon were found to be approximately 5×10^{-6} , 10, and 0.001, respectively. After investigating the optimal hyperparameters, we assumed that the same parameters were applicable to all grid cells because they did not vary extensively among the grid points. Next, the precipitation estimation performance was investigated based on the correlation coefficients of 35 grid cells, and the coefficients were averaged over every 10 grids, as shown in Figure

S3. First, the optimal gamma value was estimated using temporary values of C (10) and epsilon (0.001). Second, the optimal C value was obtained using the optimal gamma value and a temporary epsilon value. Third, the optimal epsilon value was obtained using both the optimal gamma and C values. Finally, the optimal gamma was obtained using both the optimal C and epsilon values. The parameters were considered to be optimally determined if they corresponded to the first estimates or if the correlation coefficients did not change considerably. The optimal values of gamma, C, and epsilon were approximately 5×10^{-6} , 10, and 0.001, respectively. Thus, we obtained the optimal hyperparameter values and configured them for all grid cells.

2.3 Definition of the size and resolution of the explanatory variable

Figure S4 shows plots of the correlation coefficients obtained under different explanatory variable area sizes. The performance tended to improve as the area size increased. Small area sizes of explanatory variables are believed to be insufficient for estimating precipitation with high accuracy because of the lack of information. In other words, few explanatory variables might be insufficient to explain the objective variable. Considering the performance and computational cost, we set the size of the explanatory variable to 21 by 21 grid cells in this study.

2.4 Estimation of heavy precipitation using quantile mapping

We utilized quantile mapping (Lafon et al., 2013) after applying the ML method to modify the amount of precipitation. Quantile mapping was applied by using the observed and simulated cumulative density functions of hourly precipitation. The corrected intensity of the simulated precipitation for a given quantile was determined by resampling from the observed precipitation intensity distribution with the same quantile value. Here, we performed quantile mapping using hourly observed precipitation data and ML-estimated data in January, April, July, and October for 11 years from 2007 to 2018, excluding the estimated year (Figure S2).

2.5 Prediction of local precipitation using 39-h forecasted precipitation simulations

We investigated the local precipitation-estimation performance based on the ML approach using the 39-h forecasted precipitation data from the MSM-GPV dataset, which started at 0 UTC each day in January, April, July, and October for five years (from 2014 to 2018, as 39-h forecasted data were not provided by the JMA until 2014). For this purpose, we applied the classifiers and cumulative density functions produced by the ML method in advance.

2.6 Fractions skill score

We investigated the local precipitation estimation performance using the fraction skill score (FSS) (Roberts & Lean, 2008), which is used to make spatial comparisons. The FSS indicates how the skill varies with the spatial scale. We estimated the FSS using the “verification” package in R software (Gilleland,

2015). We used the grid values in the evaluation area shown in Figure 1b to estimate the FSS by varying the box radius from 0 to 20. The FSSs were calculated by applying the observed area-averaged values of the frequency of 95th-percentile, the monthly mean, and the 95th-percentile, respectively, as the threshold for convenience.

3 Results

3.1 Validation of the bias correction method for precipitation

Model bias is clearly found in the long-term precipitation simulations. Figure 2 shows the frequency distributions of the 95th-percentile values obtained from four data sets, namely, the observed precipitation (OBS), simulated precipitation (SIM), precipitation estimated by the ML method with quantile mapping (MLQM), and precipitation estimated by quantile mapping (QM) datasets, in January and July from 2007 to 2018. Figure 2 also presents histograms of the correlation coefficient and root mean square error (RMSE) values in the evaluation area in January, April, July, and October. The frequency characteristics are estimated well by MLQM, while the QM frequencies are almost the same as those in the SIM dataset and cannot improve the biases. In addition, the frequency distributions shown in the MLQM dataset are the same as those in the ML dataset, the experiment performed before the QM method was applied. Therefore, the ML method can improve the precipitation frequency at each grid point and then modify the amount of precipitation by applying QM. The correlation coefficient and RMSE values of the frequencies in the studied region are shown in Figure 2i and 2j, respectively. The correlation coefficients in the MLQM dataset exceed 0.98, and the RMSEs are reduced significantly, while the correlation coefficient and RMSE values in the QM dataset are almost the same as those in the SIM dataset.

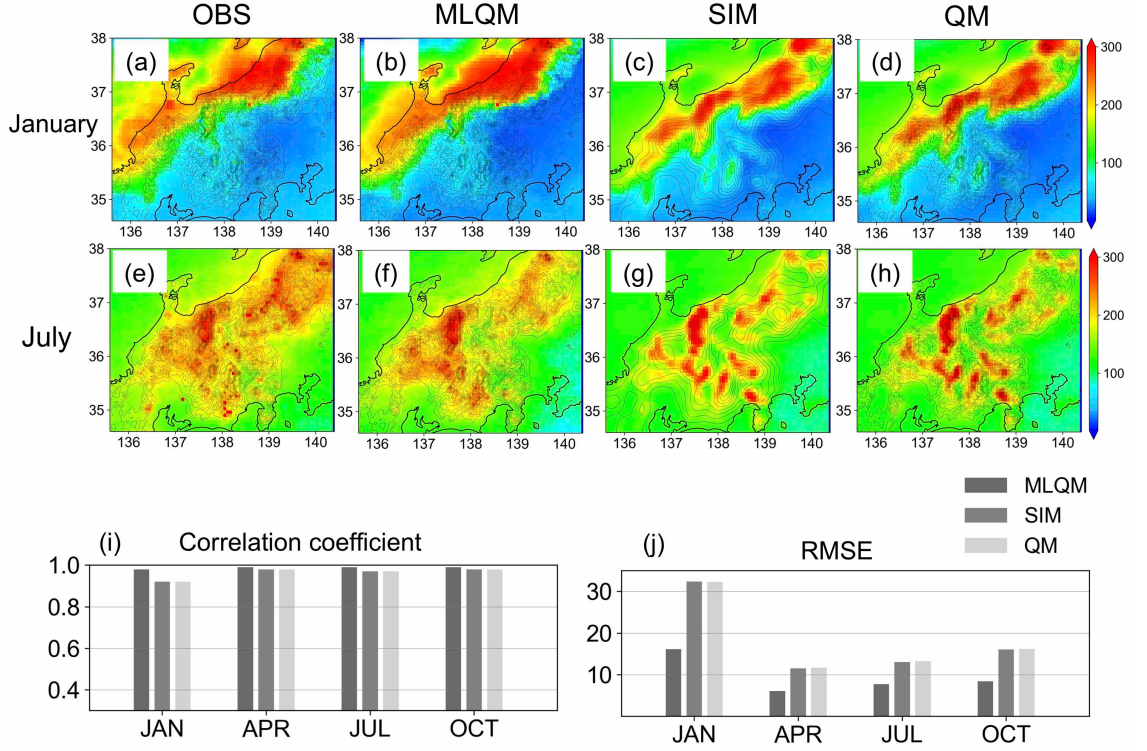


Figure 2. Frequency distributions of the 95th percentile obtained from four data sets, namely, the observed precipitation (OBS), precipitation estimated by the ML method with quantile mapping (MLQM), simulated precipitation (SIM), and precipitation estimated by quantile mapping (QM) datasets, in January and July from 2007 to 2018. The histograms show the correlation coefficient and root mean square error (RMSE) values in the evaluation area in January (JAN), April (APR), July (JUL), and October (OCT).

Figure 3 shows the distributions of the 12-year average monthly means and 95th-percentile precipitation values in January and July. QM is able to improve the monthly mean precipitation distribution characteristics to some extent, although the performance of QM is less optimal than that of MLQM. The QM method does not work sufficiently if the differences in frequency between the OBS and SIM values are significantly large because the monthly mean precipitation amount is also greatly influenced by the frequency. Unsurprisingly, QM can improve the 95th-percentile values due to the characteristics of the QM method (Maraun et al., 2017). The same features are found in the precipitation distribution characteristics in April and October (Figures S5 and S6). Figure S7 shows histograms of the correlation coefficients and RMSEs of the monthly mean and

95th-percentile values. The correlation coefficients of the monthly means in the MLQM method exceed 0.97, and the RMSEs are reduced significantly using this method, while the correlation coefficients of the monthly means in the QM method are inferior to those in the MLQM method (Figure S7a, S7b).

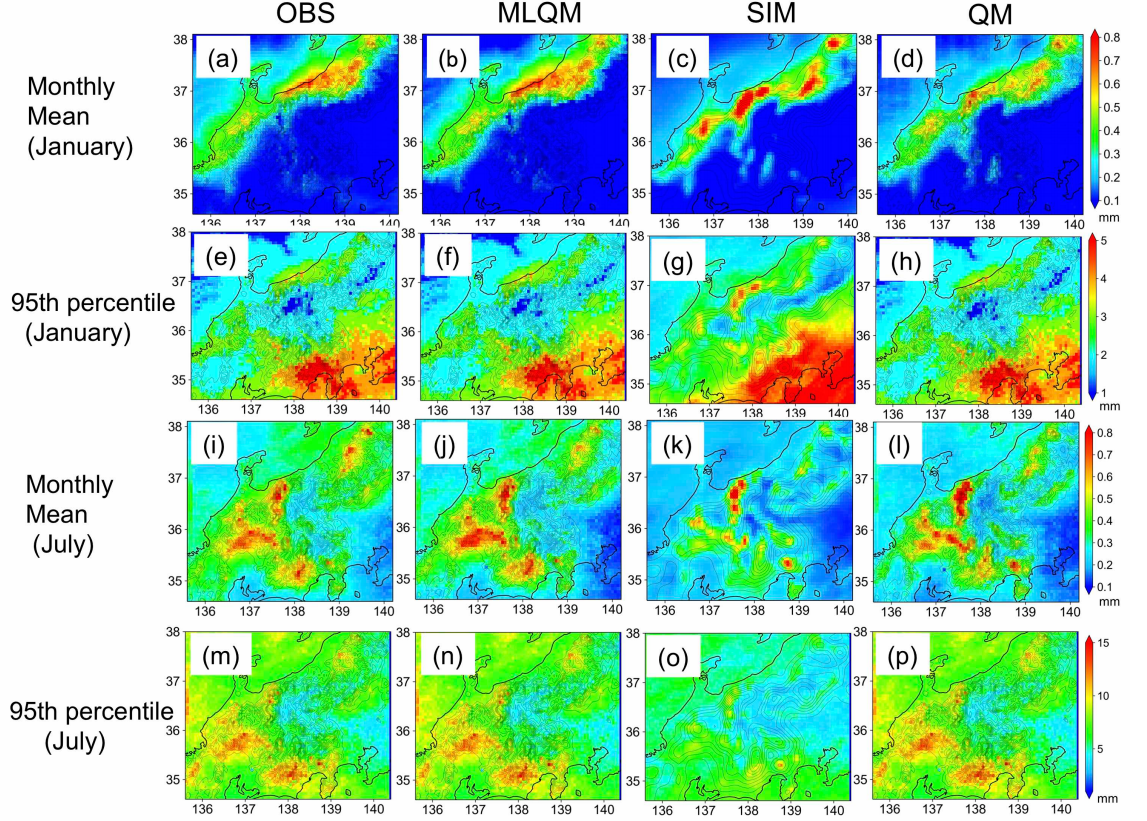


Figure 3. Distributions of the 12-year average monthly mean and 95th-percentile values in the OBS dataset and precipitation estimated by the MLQM, SIM, and QM methods in January and July.

Figure S8 shows the FSSs of the frequencies of the 95th-percentile values and the monthly mean precipitation values, and the 95th-percentile values. The FSSs of MLQM are higher than those of SIM in all months. The FSSs of the frequencies in the QM dataset are almost the same as those of SIM. It is confirmed that MLQM can modify the bias of the spatial distribution of precipitation and drastically improve the performance.

3.2 Applicability of the method for forecasting precipitation data

Figure 4 shows the frequencies of the 95th-percentile OBS data, the precipitation estimated by the ML method with QM (MLQM-Forecast), and the simulated

precipitation (SIM-Forecast) in the five years from 2014 to 2018 in January and July as well as the correlation coefficient and RMSE values in January, April, July and October. We used the forecasted data from 25 to 39 hours following the initial time to remove the influence of the initial time as much as possible. However, accurately comparing the simulated precipitation amount with the observed amount remains difficult because the forecasted data include errors that depend on initial conditions. Nevertheless, the frequency distribution characteristics are accurately estimated by MLQM-Forecast (Figure 4 and Figure S9). The correlation coefficient and RMSE values of the frequencies in the MLQM-Forecast dataset are improved in all months.

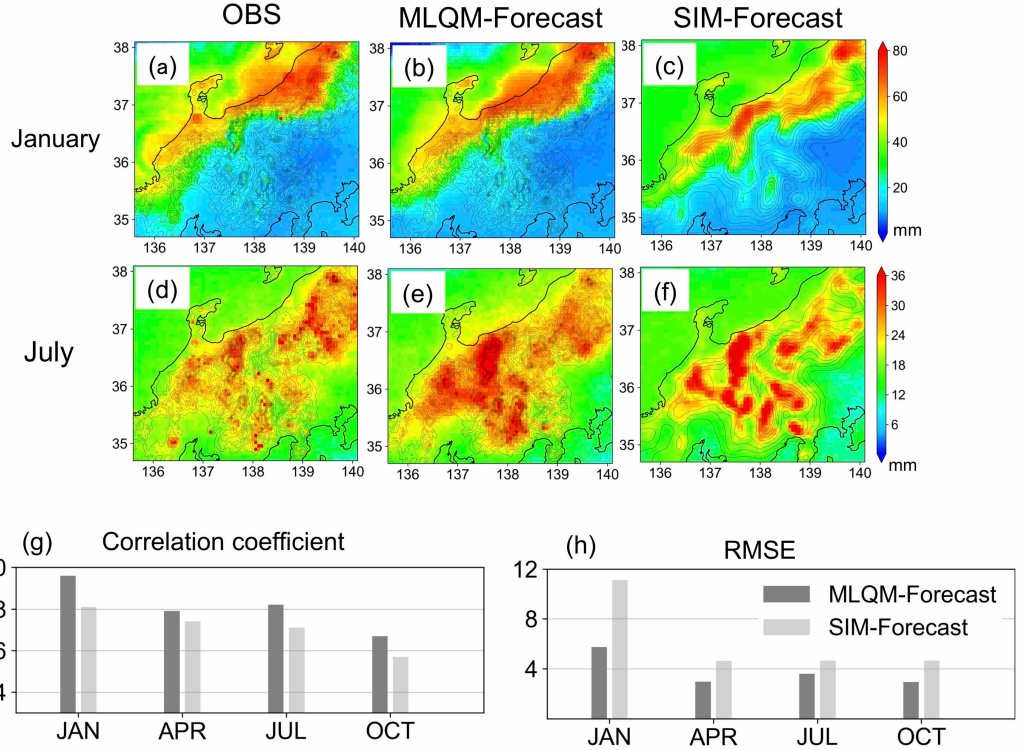


Figure 4. Distributions of the frequency of the 95th-percentile values obtained from the OBS, MLQM-Forecast, and SIM-Forecast datasets. The histograms show the correlation coefficients and RMSEs of the 95th-percentile frequencies in January, April, July, and October.

Moreover, the method reasonably improves the characteristics of the distributions of the monthly mean precipitation and 95th-percentile values in January in MLQM-Forecast (Figure S10), although the amount of precipitation is over-

estimated in both months. The same features are found in April and October (Figure S11). As a result, the correlation coefficient and RMSE values of the monthly mean and 95th-percentile values of MLQM-Forecast are worse than those of SIM-Forecast (Figure S12). The testing term (five years) might be too short to accurately evaluate the amount of precipitation because the precipitation distributions could be greatly influenced by only a few disturbances, such as intensified rain bands and typhoons, in the evaluation area during the warm season. The same characteristics are confirmed by the FSS. The FSSs of MLQM-Forecast are higher than those of SIM-Forecast in all months (Figure S13).

4 Discussion

In general, it is quite difficult to model extreme events due to the sparsity of sampled data, particularly in summer. We cannot confirm that the cumulative density functions of the simulated precipitation correspond to those of the observed precipitation at higher precipitation intensities because the return periods of the precipitation intensities are unknown. Hence, more sampling is required to model the precipitation amount accurately.

Nevertheless, the ML method has the ability to improve the estimated precipitation frequency at each grid point, as we expected. Generally, it is impossible to estimate precipitation without ML because it is too complicated to specify the ML recognition result. Even a slight change in the atmospheric fields may cause a significant change in local precipitation. Machine learning can easily find patterns in such complex weather phenomena and allows humans to recognize that there some patterns exist even in extremely complex phenomena. We can at least understand that ML reveals some sort of relation between the simulated precipitation and observed precipitation, corresponding to large-scale weather conditions, and estimates the precipitation using this relation.

The area size of the explanatory variable is also important for achieving a high modeling performance. The performance is reduced if the size is excessively small because there is insufficient information to recognize the relation of the simulations with the observations. Moreover, in this study, the simulation performance was not improved even when the size was extended to an area of more than 113 square kilometers (Figure S4). We speculate that the ML method recognizes the meso-beta scale of convection systems such as cold and warm fronts formed in low-pressure systems and the precipitation distribution characteristics corresponding to orography.

Furthermore, high-quality input observation data and simulation data are also necessary to improve the forecasting performance. The high performance in this case indicates the high quality of the observed and simulation data provided by the JMA. Moreover, it is extremely difficult to simulate convective systems at the meso-beta scale perfectly using numerical models because of the nonlinearity of precipitation systems. Therefore, subtle deviations in the location of the simulated precipitation band may affect both the learning and inference processes

of machine learning, leading to larger errors in some cases. Regarding the issue of the input data, it is necessary to include the relationships of the input data with the simulations and meteorological theories in the discussion. When using a simple input dataset, it is easier to interpret the results and find problems in the input data, even if machine learning is a black box.

In recent years, deep learning approaches have been developed rapidly, and thus, bias correction methods have become more complicated. However, while the accuracy of these methods might be improved by using complicated techniques, it is much more difficult to understand what the ML method recognizes under these conditions, and the methods might lack versatility. Moreover, when using complicated techniques, it is impossible to fix issues that arise because we cannot diagnose what is wrong with the method. Accordingly, the accuracy is not the only matter of concern; rather, it is necessary to understand the system (even if only slightly) to improve the reliability of the method. We can solve the abovementioned problem by identifying the causes of issues and steadily improving the method.

In this study, we assumed some relations between the observed and simulated precipitation values and developed an ML method to correct the precipitation distribution characteristics to reflect local conditions based on the observations and the reproducibility of the numerical model by using ML with simple input data. The hypothesis-verification approach is very important for clarifying the system as a scientific process and to begin to reduce the black-box issue of ML. We expect that this method will be used to further improve the forecasting performance by developing a combination of machine learning and quantile mapping and that this method will continue to be a powerful tool for clarifying complex weather and climate systems.

5 Conclusions

We developed a machine learning bias correction method using simple input data and verified the relations between the simulated distribution of precipitation in an area spanning approximately 100 km² and the observed precipitation at the center of the area. The method can represent the precipitation frequency at each grid point corresponding to various weather conditions, which are greatly influenced by local conditions. The combined method with QM can sufficiently modify the monthly precipitation amount. Furthermore, this simple method has many advantages: 1) it is simple to handle and process the data using this method; 2) preprocessing, such as dimension reduction, is unnecessary for the input data; and 3) the behavior of ML is relatively comprehensible, and thus, the reliability of this method can be improved. Therefore, we expect that this method can easily be verified in other regions and seasons and will be useful for estimating the water resources and risks of water disasters.

Acknowledgments and Data

This study was supported by the water environment and resource research project at the Earth Observation Research Center, Japan Aerospace Explo-

ration Agency (JX-PSPC-524437); the Environment Research and Technology Development Fund S-20 (JPMEERF21S12020) of the Environmental Restoration and Conservation Agency of Japan; the Cross-ministerial Strategic Innovation Promotion Program (SIP); the development of an application toward water-related problems in the Data Integration & Analysis System (DIAS) (JPMXD0716808979) from the MEXT, Japan; and the Integrated Research Program for Advancing Climate Models (TOUGOU) (JPMXD0717935457) from the MEXT, Japan.

The MSM-GPV data were obtained from the archive system (<http://apps.diasjp.net/gpv>) of the Kitsuregawa laboratories, respectively, at the Institute of Industrial Science, University of Tokyo. The Radar-AMeDAS data, which were provided by the JMA, are available by contacting the Japan Meteorological Business Support Center (<http://www.jmbc.or.jp/en/contact.html>) by email the International Service Officer (jmbc-ab@jmbc.or.jp). The official name of Radar-AMeDAS dataset in Japan is “Kaiseki-Uryou”.

References

- Al-Anazi, A. F., & Gates, I. D. (2012). Support vector regression to predict porosity and permeability: Effect of sample size. *Computers & geosciences*, *39*, 64-76.
- Anguita, D., Ghio, A., Greco, N., Oneto, L., & Ridella, S., (2010). Model selection for support vector machines: Advantages and disadvantages of the Machine Learning Theory. *The 2010 International Joint Conference on Neural Networks (IJCNN), Barcelona, Spain*, 1-8. doi: 10.1109/IJCNN.2010.5596450
- Bergstra, J., & Bengio, Y. (2012). Random search for hyper-parameter optimization. *Journal of machine learning research*, *13*(2), 281-305. <https://www.jmlr.org/papers/volume13/bergstra12a/bergstra12a>
- Castelvecchi, D. (2016). Can we open the black box of AI?. *Nature News* *538.7623*, 20. doi:10.1038/538020a
- Chen, H., Chandrasekar, V., Cifelli, R., & Xie, P. (2019). A machine learning system for precipitation estimation using satellite and ground radar network observations. *IEEE Transactions on Geoscience and Remote Sensing*, *58*(2), 982-994.
- Cherkassky, V., & Ma, Y. (2004). Practical selection of SVM parameters and noise estimation for SVM regression. *Neural networks*, *17*(1), 113-126. [https://doi.org/10.1016/S0893-6080\(03\)00169-2](https://doi.org/10.1016/S0893-6080(03)00169-2)
- Fan, J., Wang, X., Wu, L., Zhou, H., Zhang, F., Yu, X., ... & Xiang, Y. (2018). Comparison of Support Vector Machine and Extreme Gradient Boosting for predicting daily global solar radiation using temperature and precipitation in humid subtropical climates: A case study in China. *Energy conversion and management*, *164*, 102-111. <https://doi.org/10.1016/j.enconman.2018.02.087>

- Gilleland, M. E. (2015). Package ‘verification’. <https://cran.r-project.org/web/packages/verification/verification> (last access: 1 June 2021).
- Hall, M. A., & Lloyd, A. S. (1999). Feature selection for machine learning: comparing a correlation-based filter approach to the wrapper. *Proceedings of the Twelfth International FLAIRS Conference*, 1999, <https://www.aaai.org/Papers/FLAIRS/1999/FLAIRS99-042.pdf>
- Ishikawa, Y., & Koizumi, K. (2002). Meso-scale Analysis. *Outline of the Operational Numerical Weather Prediction at the Japan Meteorological Agency*, 26-31.
- JMA, 2019: NWP Application Products. https://www.jma.go.jp/jma/jma-eng/jma-center/nwp/outline2019-nwp/pdf/outline2019_04.pdf (accessed on 11 March 2021)
- Michelangeli, P. A., Vrac, M., & Loukos, H. (2009). Probabilistic downscaling approaches: Application to wind cumulative distribution functions. *Geophysical Research Letters*, 36, 11. <https://doi.org/10.1029/2009GL038401>
- Lafon, T., Dadson, S., Buys, G., & Prudhomme, C. (2013). Bias correction of daily precipitation simulated by a regional climate model: a comparison of methods. *International Journal of Climatology*, 33(6), 1367-1381. <https://doi.org/10.1002/joc.3518>
- LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *nature*, 521(7553), 436-444. <https://doi.org/10.1038/nature14539>
- Liu, P., Choo, K. K. R., Wang, L., & Huang, F. (2017). SVM or deep learning? A comparative study on remote sensing image classification. *Soft Computing*, 21(23), 7053-7065. <https://doi.org/10.1007/s00500-016-2247-2>
- Makihara, Y., Uekiyo, N., A. Tabata, A., & Abe, Y. (1996) Accuracy of radar-AMeDAS precipitation. *IEICE Transactions on Communications*, 79, 751-762.
- Maraun, D., Shepherd, T. G., Widmann, M., Zappa, G., Walton, D., Gutiérrez, J. M., ... & Mearns, L. O. (2017). Towards process-informed bias correction of climate change simulations. *Nature Climate Change*, 7(11), 764-773. <https://doi.org/10.1038/nclimate3418>
- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., et al. (2011). Scikit-learn: Machine learning in Python. *The Journal of machine Learning research*, 12, 2825-2830.
- Rudin, C. (2019). Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nature Machine Intelligence*, 1.5, 206-215. <https://doi.org/10.1038/s42256-019-0048-x>
- Roberts, N. M., & Lean, H. W. (2008). Scale-selective verification of rainfall accumulations from high-resolution forecasts of convective events. *Monthly Weather Review*, 136, 78-97. <https://doi.org/10.1175/2007MWR2123.1>

- Sachindra, D. A., Ahmed, K., Rashid, M. M., Shahid, S., & Perera, B. J. C. (2018). Statistical downscaling of precipitation using machine learning techniques. *Atmospheric research*, 212, 240-258. <https://doi.org/10.1016/j.atmosres.2018.05.022>
- Saito, K., Fujita, T., Yamada Y., Ishida J., Kumagai Y., Aranami K., et al. (2006). The operational JMA non-hydrostatic mesoscale model. *Mon. Wea. Rev.*, 134, 1266-1298. <https://doi.org/10.1175/MWR3120.1>
- Sivapragasam, C., Liong, S. Y., & Pasha, M. F. K. (2001). Rainfall and runoff forecasting with SSA-SVM approach. *Journal of Hydroinformatics*, 3(3), 141-152. <https://doi.org/10.2166/hydro.2001.0014>
- Smets, K., Verdonk, B., & Jordaan, E. M. (2007). Evaluation of performance measures for SVR hyperparameter selection. *In 2007 International Joint Conference on Neural Networks. IEEE*, 637-642. <https://doi.org/10.1109/IJCNN.2007.4371031>
- Smola, A. J., & Schölkopf, B., (2004) A tutorial on support vector regression. *Statistics and computing*, 14, 199-222.