

Benchmarking the geographic generalization of deep learning models for precipitation downscaling

Paula Harder, Luca Schmidt, Francis Pelletier, Nicole Ludwig, Matthew Chantry, Christian Lessig, Alex Hernandez-Garcia, David Rolnick

Angaben zur Veröffentlichung / Publication details:

Harder, Paula, Luca Schmidt, Francis Pelletier, Nicole Ludwig, Matthew Chantry, Christian Lessig, Alex Hernandez-Garcia, and David Rolnick. 2026. "Benchmarking the geographic generalization of deep learning models for precipitation downscaling." *Scientific Reports* 16 (1): 3733.
<https://doi.org/10.1038/s41598-025-34557-4>.

Nutzungsbedingungen / Terms of use:

CC BY 4.0





OPEN Benchmarking the geographic generalization of deep learning models for precipitation downscaling

Paula Harder^{1,2,6}, Luca Schmidt^{1,3,6}✉, Francis Pelletier¹, Nicole Ludwig³, Matthew Chantry², Christian Lessig², Alex Hernandez-Garcia^{1,5} & David Rolnick^{1,4}

Earth System Models (ESM) are our main tool for projecting the impacts of climate change. However, running these models at sufficient resolution for local-scale risk-assessments is not computationally feasible. Deep learning-based super-resolution models offer a promising solution to downscale ESM outputs to higher resolutions by learning from data. Yet, due to regional variations in climatic processes, these models typically require retraining for each geographical area—demanding high-resolution observational data, which is unevenly available across the globe. This highlights the need to assess how well these models generalize across geographic regions. To address this, we introduce RainShift, a dataset and benchmark for evaluating downscaling under geographic distribution shifts. We evaluate state-of-the-art downscaling approaches including GANs and diffusion models in generalizing across data gaps between the Global North and Global South. Our findings reveal substantial performance drops in out-of-distribution regions, depending on model and geographic area. While expanding the training domain generally improves generalization, it is insufficient to overcome shifts between geographically distinct regions. We show that addressing these shifts through, for example, domain adaptation can improve spatial generalization. Our work advances the global applicability of downscaling methods and represents a step toward reducing inequities in access to high-resolution climate information.

High-resolution climate projections are crucial for planning effective responses to extreme weather and climate events. With ongoing climate change, extreme events such as floods, droughts, and heatwaves are expected to become more frequent and severe, threatening infrastructure, agriculture, energy systems, and public health¹. Among these, precipitation extremes are particularly destructive, causing floods, landslides, and soil erosion; and are projected to intensify at local scales^{2,3}. Yet, accurately modeling precipitation extremes remains difficult due to their high spatial and temporal variability and the non-linear, multi-scale processes involved^{4,5}.

Earth System Models (ESMs) are the primary tools for understanding the Earth's climate system and projecting future conditions. By representing physical, chemical, and biological processes, they simulate interactions between climate, the carbon cycle, ecosystems, and human activities. However, the spatial resolution of ESMs—typically around 100 km—is too coarse to resolve small-scale processes. Instead, sub-grid processes are approximated through parameterizations that estimate their average influence. A critical example is deep convection, a major driver of precipitation and a major cause of extreme rainfall, such as flash floods and landslides⁶. As these processes occur at spatial scales finer than those resolved by ESMs, the models are unable to accurately represent these events⁴.

To address this limitation, downscaling methods are used to increase the spatial resolution of ESM output. There are two families of approaches for downscaling: dynamical downscaling and statistical downscaling. Dynamical downscaling uses a high-resolution regional climate model driven by ESM boundary conditions to simulate fine-scale processes within a limited area of interest. In contrast, statistical downscaling techniques learn empirical relationships between large-scale predictors and local-scale observations from training data and apply these learned relationships to predict localized climate outputs. Compared to dynamical downscaling

¹Mila - Quebec AI Institute, Montreal, Canada. ²European Centre for Medium Range Weather Forecasts (ECMWF), Bonn, Germany. ³Cluster of Excellence Machine Learning, University of Tübingen, Tübingen, Germany. ⁴McGill University, Montreal, Canada. ⁵Université de Montréal, Montreal, Canada. ⁶Paula Harder and Luca Schmidt contributed equally to this work. ✉email: luca-marie.schmidt@mila.quebec

methods, statistical downscaling methods are more computationally efficient but require high-resolution data for training.

Recently, deep learning methods have shown promise for statistical downscaling, leveraging advances in computer vision—particularly super-resolution techniques. Early approaches used convolutional neural networks (CNNs) to learn deterministic mappings from coarse to high resolution⁷, with architectures such as super-resolution CNNs, U-Nets^{8,9}, and ResNets^{10,11} being widely applied to climate downscaling. More recent work has shifted toward probabilistic frameworks, leveraging generative models to represent the inherent uncertainty in the downscaling task. Generative adversarial networks (GANs)¹², especially conditional generative adversarial networks (cGANs)¹² and stabilized variants like Wasserstein GANs^{13,14} are popular choices for downscaling. More recently, diffusion models have also shown strong performance in modeling complex, high-dimensional distributions^{15–19}. Among all climate variables, precipitation is particularly challenging to model and forecast due to its stochastic, high-frequency, spatial and temporal variability. Capturing these fine-scale variations and inherent uncertainties would make generative approaches successful for precipitation downscaling.

Unlike dynamical downscaling, statistical downscaling algorithms are not restricted to specific data sources or geographic regions. By construction, statistical downscaling relies on statistical relationships learned during the training phase, allowing it to be applied to new regions at inference time. However, this transferability hinges on the assumption that both predictor and predictand distributions, as well as their statistical relationships, remain stationary across space. In practice, this assumption does not hold due to substantial geographic variability in topography, climatic conditions and processes. For example, precipitation in equatorial regions is strongly driven by convection and, therefore, substantively different than in Europe and North America. Consequently, statistical downscaling models often require retraining for each target region, which relies on the availability of high-resolution observational data.

Despite the large amount of weather and climate datasets, high-quality observations are unevenly distributed globally. Ground-based radar and gauge data, essential for training and validating downscaling models, are particularly sparse in many parts of the Global South (see Fig. 1). Yet, it is also these regions that are often the most exposed and vulnerable to climate change and extreme weather events like heavy rainfall and flooding³. This global imbalance in data availability—coupled with heterogeneous regional climate processes—presents challenges in generalization of deep learning-based downscaling models.

To address these challenges, we introduce RainShift, a large-scale global benchmark and dataset designed to evaluate the geographical generalization of deep learning-based downscaling. RainShift defines test scenarios where models are trained on subsets of data-rich regions and tested on regions with scarce availability of high-resolution observations. The objective is to support the development of models that can be applied to regions with little or no high-resolution data, i.e. considering zero-shot or few-shot settings. In practice, we select evaluation regions that serve as proxies for such data-scarce regions, allowing us to evaluate and benchmark model performance. The dataset is built from ERA5 reanalysis and IMERG satellite precipitation data. We establish baseline results by evaluating state-of-the-art models for probabilistic precipitation downscaling, including GANs and diffusion-based architectures. RainShift is intended to support the development of



Fig. 1. Map of ground-based radar stations. The map shows the availability of precipitation data, with each blue dot representing a station. Coverage is relatively high in the Global North and comparatively low across the Global South. Image from the Tropical Globe radar database²⁰.

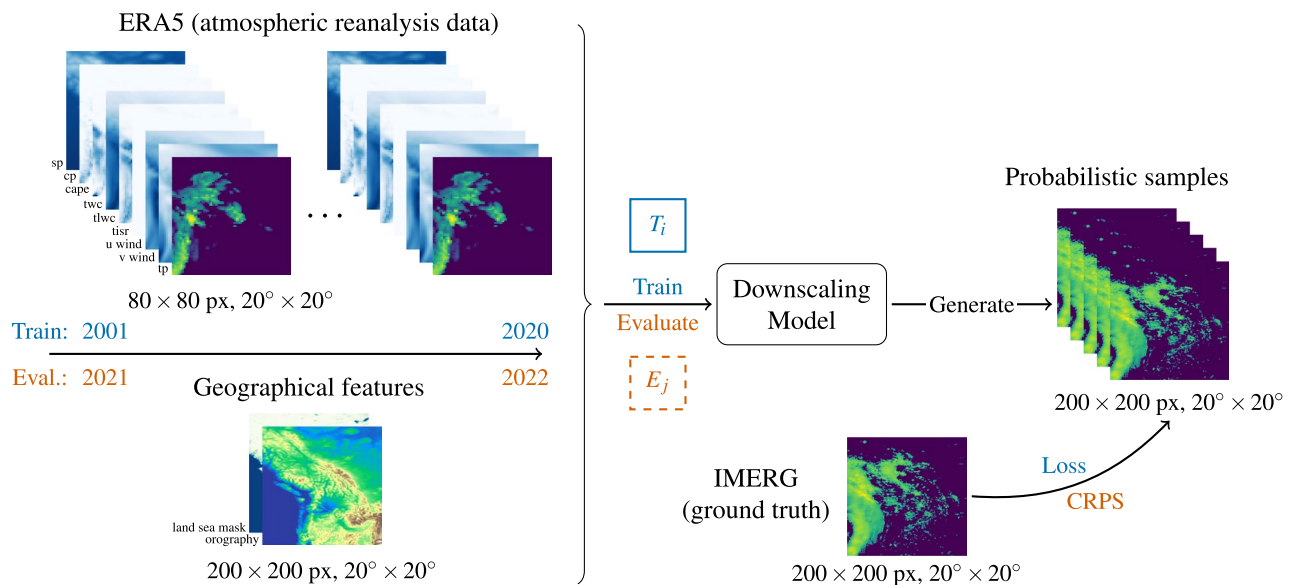


Fig. 2. Graphical summary of RainShift setup. The inputs of the downscaling model are a combination of ERA5 time series data and geographical features. The downscaling model is then able to generate probabilistic samples. For training, we sample from geographic areas $T_1, \dots, T_{12}$ and years 2001–2020, and compare the generated samples with the ground truth target IMERG to compute the loss. For evaluation, we use areas $E_1, \dots, E_6$, and years 2021–2022.

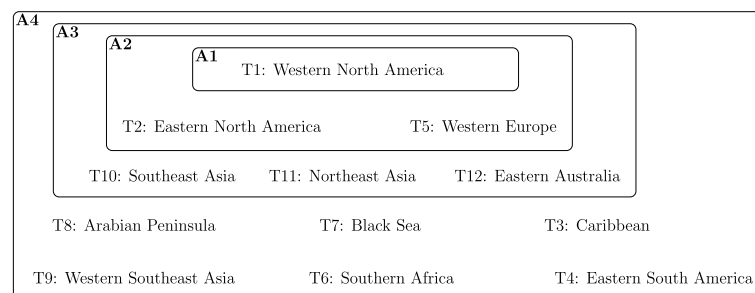


Fig. 3. Illustration of training configurations. The training configurations A1, ..., A4 are composed of progressively larger subsets of the 12 selected training regions located in the Global North. The choice of regions is guided by availability of high-resolution observational data and inspired by existing works^{14,23}.

approaches that generalize to low-data regions, particularly in underrepresented areas such as the Global South. Our main contributions are the following:

- We frame the task of downscaling across a geographically varying data distribution, based on a critical gap in Earth system modeling.
- We introduce the RainShift benchmark dataset, along with tools for expanding it to new regions, data loaders, training pipelines and evaluation frameworks, within an accessible library to facilitate further research.
- We evaluate a variety of state-of-the-art machine learning downscaling models on RainShift and find substantial variation in spatial generalization across models and regions; we show that domain adaptation through data alignment and fine-tuning techniques can improve performance especially in cases where generalization is limited by strong geographic distribution shifts.

Results

We train a variety of downscaling models in scenarios representing data-abundant regions and then evaluate their performance in data-sparse regions, mainly in the Global South. The downscaling task consists of learning a mapping from coarser resolution reanalysis data (ERA5) to higher resolution satellite-based precipitation (IMERG), as illustrated in Fig. 2. These globally consistent data sources enable a controlled benchmark setup. To evaluate spatial generalization, we define 12 training regions and 6 evaluation regions. We combine the training regions into four progressively larger training scenarios (A1, ..., A4) to simulate varying levels of observational coverage (see Fig. 3); their geographic distribution, along with that of the evaluation regions, is shown in Fig. 4.

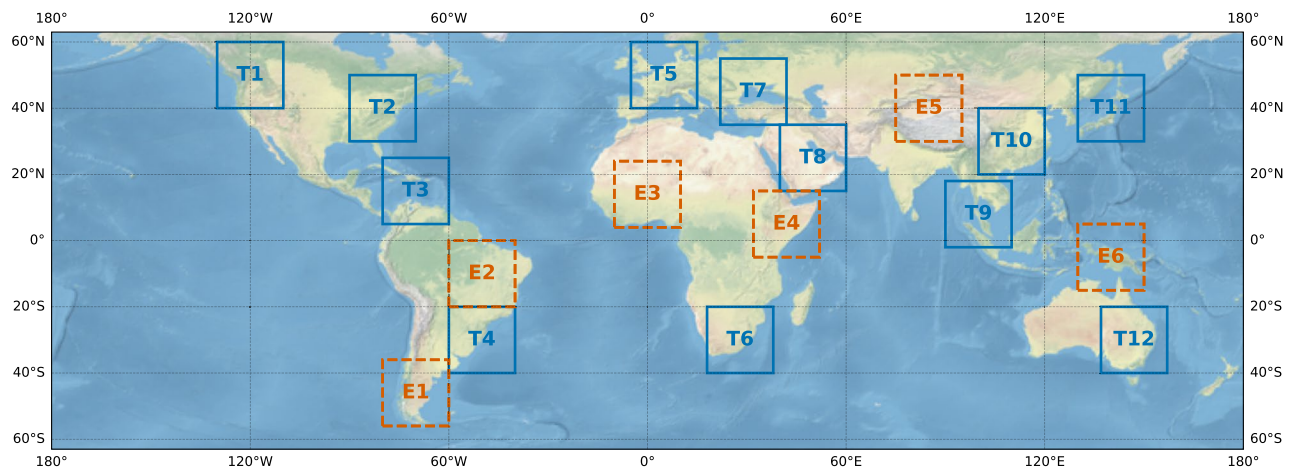


Fig. 4. Illustration of location splits and training configurations. Patches $T_1, \dots, T_{12}$ represent training regions and patches $E_1, \dots, E_6$ correspond to evaluation areas that are used within 4 sub-tasks, simulating different scenarios that correspond to varying levels of data availability.

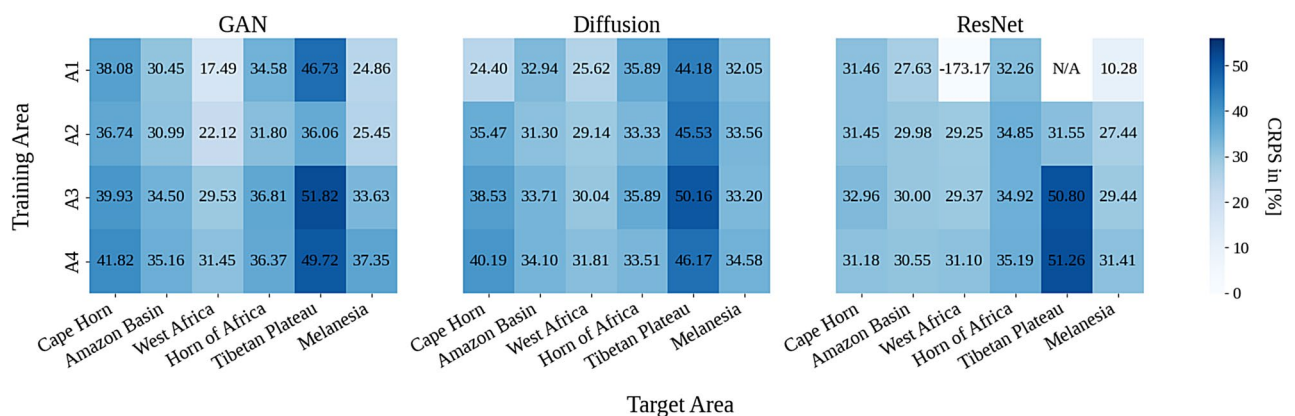


Fig. 5. Heatmap of % improvement relative to interpolation. Change in CRPS (lower better) in [%] for each model relative to bilinear interpolation. $A_1, \dots, A_4$ represent hierarchical training scenarios with progressively more high-resolution data, from training on a single region (A_1) to using several training regions across the Global North (A_4). Positive values show improvements over the interpolation baseline. Model performance generally improves when expanding the training areas from A_1 to A_4 , but this trend depends strongly on the geographical region. The N/A value indicates an instance where the CRPS value is not reliable due to numerical instabilities during inference.

The analysis includes several state-of-the-art downscaling approaches: ResNets¹³ (a deterministic approach), Wasserstein GAN with gradient penalty²¹ (a probabilistic approach), and a diffusion-based method²² (also probabilistic). We also include a simple bilinear interpolation as a baseline. We evaluate model performance with respect to spatial generalization using the continuous ranked probability score (CRPS) and the percentile Mean Absolute Error (MAE), considering out-of-distribution performance both in absolute terms and relative to in-distribution evaluation. CRPS measures the performance of probabilistic models by quantifying the average error of a stochastic prediction with respect to a reference value. To evaluate extremes, we use the percentile MAE, which measures errors in the upper percentiles of the prediction and reference. Our analysis investigates key factors affecting generalization, including the downscaling model, the geographic region, and size of the training domain.

To address the observed performance drops in unseen regions, we propose two domain adaptation strategies. For the zero-shot setting, we assume that no high-resolution target data is available, and therefore propose a simple data alignment approach based on quantile mapping. In the few-shot setting, where a small amount of labeled target data is available, we fine-tune the pretrained models to improve their spatial generalization ability.

ML-based downscaling adds value over simple interpolation

All learned models demonstrate some generalization to unseen target regions, as shown by their consistent improvement over the bilinear interpolation baseline (see Fig. 5). These improvements are particularly

pronounced for the percentile MAE (see Fig. 6), with relative improvement of up to 83%. This suggests that the added value of downscaling is especially relevant for capturing extreme events. The magnitude of these improvements, however, differs strongly across regions. For CRPS, relative reductions range between roughly 30% and 50%. Bilinear interpolation performs worst in Melanesia, which shows the highest baseline errors, whereas the Tibetan Plateau represents the easiest out-of-distribution target region with the lowest baseline errors. Interestingly, although West Africa and the Horn of Africa have similar mean precipitation, interpolation performs worse in West Africa—likely due to stronger mismatches between ERA5 and IMERG or greater precipitation variability. Overall, the results highlight the general validity and robustness of the ML-based downscaling approach, indicating that such models can be used to downscale coarse-resolution inputs even when applied to regions not seen during training.

Probabilistic models outperform deterministic ones

Across almost all evaluation regions, probabilistic generative approaches, both GAN and diffusion models, consistently outperform the deterministic ResNet model. While the differences between models are less pronounced for CRPS (see Table 1), substantial gaps appear in terms of percentile MAE in Fig. 6 and Supplementary Fig. S1–S2 and Tables S3–S4. In Fig. 6, both probabilistic models show clear improvements over bilinear interpolation, while ResNet shows only limited improvements. In addition, ResNet even shows unstable behavior under strong distribution shifts, e.g. in region E5 (Tibetan Plateau) under configuration A1, where inference-time instabilities make CRPS values unreliable. These findings highlight that probabilistic models are more robust to geographic distribution shifts and provide greater added value, particularly for extreme events.

Diffusion is more robust under extremes and limited data

GANs and diffusion models perform comparably in terms of CRPS. In the largest training setup (A4), they reach similar absolute improvements (Table 1) and relative gains (Fig. 5) over the interpolation baseline. However, the two models differ in how they respond to extending the training domain. The diffusion model already generalizes well when trained on a smaller training domain (A1), whereas the GAN shows more gradual improvements with increases in training data. This suggests that diffusion models are more robust to limited training data, while GANs benefit more from larger, more diverse training datasets.

The differences between the two models becomes greater for extreme-event metrics. Diffusion mostly outperforms GANs, with the clearest advantage in Amazon Basin and Melanesia—regions characterized by the highest precipitation intensities, strongest variability, and largest ERA5–IMERG mismatches in the extremes (see Supplementary Table S1). In Cape Horn, a drier region with fewer extremes, diffusion and GANs perform nearly equally, reflecting the limited advantage of extreme modeling.

Expanding the training area is not sufficient to overcome generalization drops

Model performance generally improves when expanding the training areas from A1 to A4 (see e.g., Fig. 5 for CRPS and Fig. 6 for percentile MAE). In some regions—such as Cape Horn, Amazon Basin, West Africa and Melanesia—expanding the training domain leads to clear improvements in CRPS and extreme-event representation. However, this trend is less consistent in other regions. In the Horn of Africa and the Tibetan Plateau, adding more training data does not necessarily lead to better generalization ability. Furthermore, the benefits of expanding training area tend to diminish with larger training domains. For the ResNet model, performance on extremes does not improve with additional training data. This suggests that increasing training data alone does not guarantee continued improvements in spatial generalization.

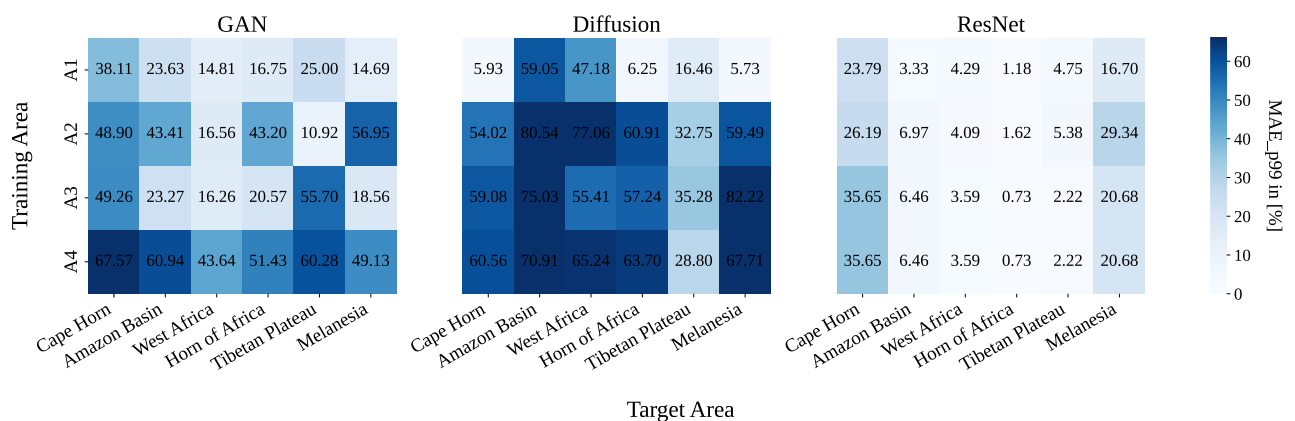


Fig. 6. Heatmap of % performance drop between in and out-of-distribution training. Change in CRPS (lower better) in [%] for each model relative to training the model directly on the target regions. A₁, ..., A₄ represent hierarchical training scenarios with progressively more high-resolution data, from training on a single region (A₁) to using several training regions across the Global North (A₄). Many regions show large negative values, indicating substantial drop in performance when evaluating on regions out-of-distribution, highlighting challenges in spatial generalization. The N/A value indicates an instance where the CRPS value is not reliable due to numerical instabilities during inference.

Model	Training Areas	Target Areas					
		E1	E2	E3	E4	E5	E6
		Cape Horn	Amazon Basin	West Africa	Horn of Africa	Tibetan Plateau	Melanesia
ERA5 Interp.	-	0.120	0.258	0.113	0.084	0.050	0.434
GAN	A_1	0.075	0.179	0.093	0.055	0.027	0.326
DM	A_1	0.091	0.173	0.084	0.054	0.028	0.295
ResNet	A_1	0.083	0.187	0.308	0.057	18113.558	0.389
GAN	A_2	0.076	0.178	0.088	0.057	0.032	0.324
DM	A_2	0.078	0.177	0.080	0.056	0.027	0.288
ResNet	A_2	0.083	0.181	0.080	0.055	0.034	0.315
GAN	A_3	0.072	0.169	0.080	0.053	0.024	0.288
DM	A_3	0.074	0.171	0.079	0.054	0.025	0.290
ResNet	A_3	0.081	0.181	0.080	0.055	0.025	0.306
GAN	A_4	0.070	0.167	0.077	0.054	0.025	0.272
DM	A_4	0.072	0.170	0.077	0.056	0.027	0.284
ResNet	A_4	0.083	0.179	0.078	0.055	0.024	0.298
GAN	E_i	0.077	0.180	0.072	0.053	0.025	0.272
DM	E_i	0.070	0.171	0.075	0.054	0.023	0.275
ResNet	E_i	0.078	0.181	0.078	0.057	0.025	0.311

Table 1. Accuracy of models on spatial generalization tasks. Test CRPS (lower better) for precipitation in [mm/h] on the designated evaluation area and averaged over test years 2021 – 2022. Shown is the mean, pixel-wise CRPS of 8 ensemble members. For deterministic models (ResNet and bilinear interpolation of ERA5 precipitation data), the MAE is shown. We report the mean precipitation amount in input and target data in [mm/h]. Best scores per subtask are in bold. The last three rows (E_i) are not contestants in the benchmark but show what is possible when training directly on the target.

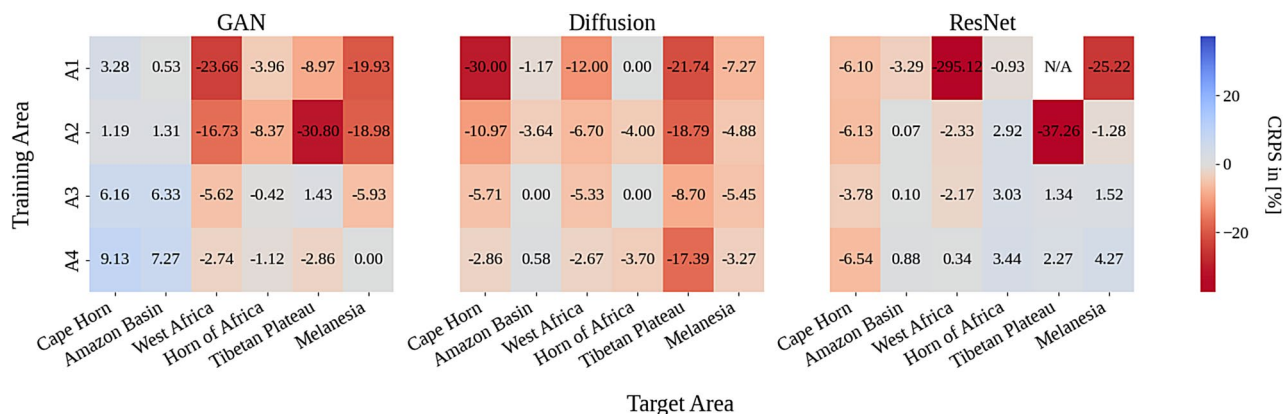


Fig. 7. Heatmap of % improvement relative to interpolation. Change in 99th percentile MAE (lower better) in [%] for each model relative to bilinear interpolation. Both probabilistic models show substantial improvements over bilinear interpolation, while the deterministic ResNet shows only limited improvement.

In-distribution training does not always lead to the best performance

For most regions and models, in-distribution training leads to the best CRPS and percentile MAE scores. Consequently, we observe substantial relative performance drops in Figs. 7 and 8 when models are trained only on data from the target region. In the Amazon Basin and Cape Horn, the size of the training domain appears more important than its overlap with the target region, with A3 and A4 performing as well or even better than in-distribution training (see rows E_i in Table 1). In the remaining four regions, however, models trained directly on the target region still perform best. The ResNet model shows more mixed results, with clear gains from on-target training in some regions, but less consistent patterns in others. Overall, these findings show that sufficiently large training areas can sometimes match the performance of on-target training, but in many cases remain insufficient depending strongly on the region and model. This highlights the importance of developing techniques that improve generalization to unseen regions.

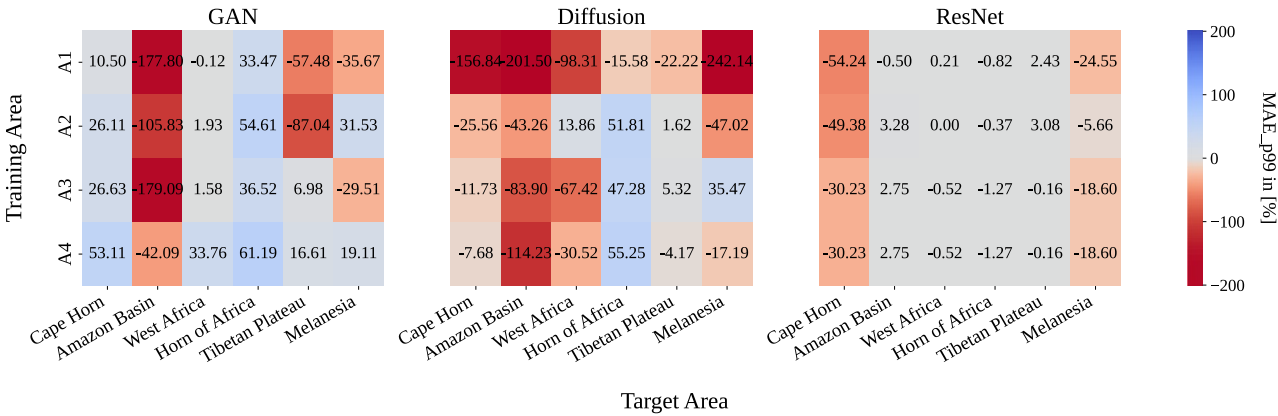


Fig. 8. Heatmap of % performance drop between in and out-of-distribution training. Change in 99th percentile MAE (lower better) in [%] for each model relative to training the model directly on the target regions.

		Target Areas					
Model	QM	E1	E2	E3	E4	E5	E6
		Cape Horn	Amazon Basin	West Africa	Horn of Africa	Tibetan Pl.	Melanesia
DM	No	0.091	0.182	0.084	0.054	0.028	0.295
DM	Yes	0.092	0.182	0.077	0.054	0.024	0.310
GAN	No	0.075	0.179	0.093	0.055	0.027	0.326
GAN	Yes	0.093	0.179	0.080	0.055	0.024	0.310
ResNet	No	0.083	0.187	0.308	0.057	18113.558	0.389
ResNet	Yes	0.093	0.181	0.079	0.055	0.025	0.310

Table 2. Impact of quantile-based domain adaptation in a zero-shot setting. CRPS (lower better) of models trained on scenario A1 (Western North America) and evaluated on the designated target regions for precipitation in [mm/h] over test years 2021 – 2022. Results show model accuracy with and without applying quantile mapping (QM) to align the input distributions between training and evaluation regions. All predictions are transformed back to the original target domain data range and re-normalized before computing the metrics. Applying quantile mapping equals or improves performance across models and most regions (except Cape Horn) demonstrating the potential of data alignment techniques to enhance spatial generalization under geographical distribution shifts.

Geographical variation dominates generalization performance

Bilinear interpolation of input precipitation reveals substantial differences in prediction accuracy between target regions, with particularly high interpolation error in the Amazon Basin (E2) and Melanesia (E6). A similar trend is reflected in the performance of deep learning models, where regions with higher precipitation remain more difficult to predict, both in absolute values (see e.g., Table 1) and in relative improvement (see e.g., Figure 5) over the interpolation baseline. This finding is consistent between different model architectures (GANs, diffusion model, and ResNet) and between different training set-ups (A1–A4), indicating that weak interpolation performance coincides with limited spatial generalization.

While generative models outperform ResNet, the performance gap between GANs and diffusion is small for CRPS (see Table 1) and nearly disappears after applying domain alignment (see e.g., Table 2). Overall, the dominant factor influencing generalization performance is the geographical region considered. Across all models, distributional shifts across regions with large climatic differences, e.g. Melanesia, are consistently more challenging to predict. We have included additional predictors e.g. orography that improve accuracy by capturing local climate conditions, but they cannot fully explain regional precipitation. This is because precipitation depends not only on conditions but also on dynamics, which vary across regions. Large-scale circulation systems, for example, can produce very different rainfall regimes even in areas with otherwise similar climate conditions. This indicates that spatial generalization is more constrained by geographic and climatic variability than by architectural choices alone, highlighting the importance of addressing such shifts through improved domain alignment or region-aware modeling strategies.

Domain adaptation techniques improve geographical generalization

In preliminary experiments, we explore two simple domain adaptation techniques within the zero-shot and the few-shot setting. In the former setting, we apply quantile mapping to align the precipitation input distribution

of the training and target regions. This is achieved by matching the cumulative distribution functions of the precipitation inputs in the target region to that of the training region. Figure 11 illustrates the distributional discrepancies, which are particularly pronounced in regions like Tibetan Plateau and Melanesia. After applying quantile correction, the CDFs are much more closely aligned, leading to improved performance across nearly all regions (see Table 2). Furthermore, we perform model fine-tuning to a small subset of the target data simulating a few-shot setting, i.e. a scenario where a limited number of such samples exists (see results in Table 3). We find that simple domain alignment or access to a small number of high-resolution samples from the target region can substantially improve model performance.

Discussion

This paper introduces the RainShift benchmark to evaluate the out-of-distribution generalization of deep learning models for precipitation downscaling across geographically distinct regions. Our dataset is constructed from paired reanalysis inputs and satellite-based precipitation targets. It covers twelve training areas with dense observational coverage and six evaluation areas located in the Global South, where high-quality observations are sparse. The benchmark task evaluates a model’s ability to learn from high-resolution training data and generalize to unseen regions, a crucial requirement to the real-world deployment of these models.

Our results show that all evaluated models add substantial value over raw input precipitation, with probabilistic generative approaches (GAN and diffusion) showing better performance than deterministic ones, both within training domains and in generalizing to unseen geographies. However, performance still largely degrades when applied to unseen regions, highlighting current limitations in cross-regional generalization. Our results indicate that spatial generalization is more strongly limited by geographic variability and distribution shifts than by differences in model architecture or training data size, indicating that addressing these shifts is a critical path forward.

To advance progress in this area, we frame cross-regional generalization as a central challenge for current downscaling models and encourage methodological innovation to address it. Our results show that simple domain alignment techniques, such as input data alignment or fine-tuning, can enhance model generalization across regions. Beyond this, other promising directions in this regard could include, for example, the use of location-aware embeddings as auxiliary inputs²⁴, or unsupervised domain adaptation methods that align feature distributions between labeled and unlabeled target domains²⁵. Further strategies such as meta-learning²⁶, which enables faster adaptation to new regions with limited data, and the integration of physical constraints to improve model robustness²⁷, also have potential to improve spatial generalization.

Our choice of training and evaluation regions is guided by the availability of high-resolution observations. To facilitate development and deployment of new methods, we use a globally homogeneous dataset, IMERG, as the target. However, a promising avenue for future work is to expand the benchmark to incorporate additional, more localized sources, such as precipitation radar data from individual countries. Models trained on such data could then be applied to regions where such high-quality observational data is truly unavailable, including much of the Global South.

The objective of RainShift is to provide a framework for evaluating the geographical generalization of downscaling models to regions most in need of accurate, high-resolution climate information. By selecting evaluation regions as proxies for truly data-scarce areas, RainShift allows testing model performance under realistic zero-shot conditions. More broadly, the RainShift benchmark is designed to bridge the gap between highly localized model development and global applicability.

Methods

We introduce the RainShift dataset, along with its preprocessing pipeline, baseline models, and evaluation framework.

RainShift dataset

The RainShift dataset builds on three global data sources: atomospheric reanalysis (ERA5), satellite-based precipitation estimates (IMERG), and two invariant geographical features—land-sea mask and orography. The use of globally consistent satellite and reanalysis data enables a controlled benchmark setup that is essential for

Model	Scenario	Target Area	CRPS	
			fine-tuning	∅ fine-tuning
ResNet	A1	Tibetan Plateau	0.031	18113.558
ResNet	A1	West Africa	0.079	0.308
ResNet	A1	Melanesia	0.309	0.389
ResNet	A2	Tibetan Plateau	0.026	0.034
GAN	A2	Tibetan Plateau	0.025	0.032

Table 3. Impact of fine-tuning model to target-region in the few-shot setting. Test CRPS (lower better) for precipitation in [mm/h] averaged over test years 2021 – 2022. Model-scenario-target area combinations where selected based on largest observed performance drops relative to in-distribution training in Fig. 7. Even a small number of high-resolution samples from the target region can substantially improve model performance.

evaluating how well downscaling models generalize spatially. Using satellite-derived precipitation as the target allows for evaluation in data-sparse regions like the Global South, while limiting confounding differences due to differences between data products.

Input data

As low-resolution input data, we use ERA5²⁸, the fifth-generation atmospheric reanalysis product of the European Center for Medium-Range Weather Forecasts (ECMWF). Reanalysis data are the result of combining historical observations with Earth system models through data assimilation to obtain global estimates of the observed climate. It provides hourly global data at $0.25^\circ \times 0.25^\circ$ resolution (approximately 25 km per pixel in mid-latitudes) on a regular latitude-longitude grid and spans years from 1950 to the present. For compatibility with IMERG, we use data from 2001 onward.

We select nine input variables based on meteorological relevance to predict subgrid rainfall variability guided by the ecPoint model²⁹ and domain-specific knowledge¹³ (see Supplementary Table S2). We additionally include geographic covariates at 0.1° resolution: (i) A land-sea mask indicating the land fraction per pixel, and (ii) an elevation map (geopotential height at surface). These predictors help capture distinct climatic conditions across regions that shape local weather.

Target data

The Integrated Multi-satellite Retrievals for GPM (IMERG)³⁰ is a product of NASA's Global Precipitation Measurement (GPM) mission and serves as high-resolution target data. IMERG provides precipitation estimates based on the GPM satellite constellation and additional observations such as gauge data. IMERG has full coverage between 60°N and 60°S at 0.1° resolution (about 10 km per pixel) on a regular latitude-longitude grid. We use the IMERG V07 Final Run product³¹, averaging its half-hourly data to hourly to match ERA5's temporal resolution. Data was accessed via NASA Goddard's GES DISC.

Summary statistics

A Table showing the average amount of precipitation, the maximum precipitation value, the standard deviation of all values, the 99th percentile precipitation, the probability of zero, as well as the wet intensity, i.e. the average precipitation rate conditional on non-zero precipitation can be found as Supplementary Table S1.

Data processing

The datasets are downloaded and preprocessed for consistency. Global data is divided into subregions and stored in the ML-friendly Zarr format³². Zarr archives retain metadata (e.g., latitude, longitude, timestamps) and are chunked for efficient loading during training. Each chunk contains 200 timesteps, optimized for sampling speed: about 10 MB per chunk for ERA5 variables and 30 MB for IMERG precipitation.

Precipitation estimates from model-based products such as ERA5 are generally considered less accurate than those from satellite-based products such as IMERG³³. To mitigate known artifacts in ERA5 precipitation data—particularly the overestimation of weak precipitation, mis-detection of non-precipitation events³⁴, and unrealistically large values—we clip the precipitation variable using the lowest and highest values from IMERG as thresholds. Very small values are set to zero, while large values are adjusted downward to align with the IMERG data distribution.

The two precipitation variables tp (total precipitation) and cp (convective precipitation) as well as orography are log-transformed by $\tilde{x} = \log(x \cdot 1000 + 1e - 5)$. The land-sea mask remains unchanged with values between 0 and 1.

Each variable is then standardized via Z-score normalization using statistics computed across all time steps, latitudes and longitudes, and aggregated across all training regions. The global mean is calculated as the average of region-wise means, and the overall variance is estimated using the pooled variance, i.e.

$$\tilde{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (\sigma_i^2 + \mu_i^2) - \left(\frac{1}{n} \sum_{i=1}^n \mu_i \right)^2,$$

where n is the number of training regions, μ_i is the mean, and σ_i^2 the variance of region T_i .

For the diffusion model, we additionally perform bilinear interpolation to the ERA5 inputs to match the spatial resolution of IMERG for compatibility with the UNet architecture. Rather than directly predicting the high-resolution target, the UNet is trained to learn the residual between the fine-resolution target and the interpolated coarse-resolution input, following prior work²².

Temporal splits

We treat every time step as an independent sample. However, as the data is a time series, we split training and evaluation data temporally. The training data in areas $T_1, \dots, T_{12}$ covers years 2001–2020, the testing data in areas $E_1, \dots, E_6$ covers years 2021–2022.

Location splits

We create RainShift choosing 18 regions worldwide, covering all continents and climate zones, each spanning $20^\circ \times 20^\circ$, as shown in Fig. 4. These regions are divided into 12 training regions ($T_1, \dots, T_{12}$) and six evaluation regions ($E_1, \dots, E_6$). The six evaluation regions are: Cape Horn, Amazon Basin, West Africa, Horn of Africa, Tibetan Plateau and Melanesia. Training regions are selected from areas with high observational data availability

(see Fig. 1 for the example of radar data), while evaluation regions are located in data-scarce areas, primarily in the Global South.

Downscaling task formulation

The RainShift benchmark focuses on a probabilistic downscaling task, where the goal is to learn the conditional distribution, $p(y|x)$, of high-resolution precipitation, y , given a low-resolution forecast and invariant features, x . A generative model, G , is trained to approximate the true distribution

$$G(x) \sim p_G(\cdot|x) \text{ such that } p_G(\cdot|x) \approx p(\cdot|x).$$

Here, the high-resolution 2D target sample $y \in \mathbb{R}^{h_h \times w_h}$ is a single-channel precipitation field from the IMERG satellite data. The input x consists of both low-resolution and high-resolution components: $x = (x_\ell, x_h)$, with $x_\ell \in \mathbb{R}^{c \times h_\ell \times w_\ell}$ representing a low-resolution c -channel forecast (here from ERA5) and $x_h \in \mathbb{R}^{d \times h_h \times w_h}$ denoting d -channel high-resolution invariant features (land-sea mask and orography). With an upsampling factor $N \in \mathbb{R}^+$, the high-resolution dimensions are given by $h_h = N \cdot h_\ell$ and $w_h = N \cdot w_\ell$. In this benchmark, $N = 2.5$ and image patches are square: $h_\ell = w_\ell = 80$ and $h_h = w_h = 200$. The channel dimensions are $c = 9$ and $d = 2$. This downscaling task is illustrated in Fig. 2.

Geographical generalization

Unlike existing benchmarks that evaluate models within the same geographic region, RainShift is designed to assess generalization across different geographies. We consider 12 different training regions and 6 regions for evaluation. Given a training area, A , the corresponding local data distribution p_A and samples $(x_A, y_A) \sim p_A$, the goal is to learn the distribution $p_E(\cdot|x_E)$ in a separate evaluation area E . Here, training and evaluation areas are disjoint, $A \cap E = \emptyset$. The task is a zero-shot prediction, with no available labels in the evaluation set. A generative model G , is desired to approximate the target distribution

$$G(x) \sim p_G(\cdot|x) \text{ such that } p_G(\cdot|x_E) \approx p_E(\cdot|x_E).$$

Training sub-tasks

To simulate different scenarios corresponding to varying levels of data availability, we define four training sub-tasks, as illustrated in Fig. 4. Each sub-task is associated with a subset of training areas $A_i \subseteq T = \bigcup_{j=1}^{12} T_j$ for $i = 1, \dots, 4$. The subsets are hierarchical—that is, $A_i \subseteq A_{i+1}$ for $i = 1, \dots, 3$ reflecting varying levels of observational availability:

1. $A_1 := T_1$: a scenario resembling the common task of just training in one geographic area.
2. $A_2 := T_1 \cup T_2 \cup T_5$: inspired by existing works that use both North American and European datasets^{14,23}.
3. $A_3 := A_2 \cup T_{10} \cup T_{11} \cup T_{12}$: an extension of A_2 adding three areas with high availability of observational data: two areas in Eastern Asia and one in Eastern Australia.
4. $A_4 := \bigcup_{i=1}^{12} T_i$: a very optimistic scenario that assumes access to a vast amount of high-resolution observations from a variety of regions and sources, including places with limited data availability.

Usage and contribution

The RainShift benchmark dataset is hosted via Hugging Face at <https://huggingface.co/datasets/RainShift/rainshift>. The dataset consists of a zipped Zarr directory per region, resulting in 300 GB of overall data. The download, data-loading, and training instructions can be found in our repository, which will be made available upon acceptance. All baselines are made available via our repository.

For the purpose of a unified benchmark task, we fix different sets of training and evaluation regions. However, we provide the tools to add new areas of interest to enable optimization or testing in any location worldwide. The instructions to create new regional subsets are provided in our repository, that will be made available upon acceptance.

Baseline models

We evaluate a diverse set of baseline models to provide a comprehensive view of performance across different model classes. These include a deterministic ResNet model, as well as two successful probabilistic methods: generative adversarial networks (GANs) and a diffusion-based approach. To contextualize these results, we also include bilinear interpolation as a simple baseline to establish a lower performance bound, and in-region training as an upper bound on expected performance.

Interpolation baseline

As precipitation is both an input and output variable, we may construct a simple baseline by bilinearly interpolating ERA5 total precipitation to approximate IMERG data. This helps assess whether poor model performance in specific regions stems from generalization issues or challenges inherent to the input data.

ResNet

Super-resolution convolutional neural networks (SRCNNs) were the first deep learning models applied to downscaling^{7,35}. These early approaches were extended with deeper network architectures and by incorporating auxiliary predictors such as topography^{35,36}. Since then, more advanced architectures have been introduced for climate downscaling. Among these, UNets³⁶ and residual networks (ResNets)^{37,38} are among the most commonly

used architectures for precipitation downscaling. In the deterministic downscaling setting, ResNet architectures still achieve state-of-the-art performance^{27,39}. Motivated by this, we choose a fully-convolutional ResNet as our deterministic baseline. The model architecture follows prior work that has demonstrated strong performance in precipitation downscaling¹³.

Generative Adversarial Networks

GANs are a popular approach for probabilistic downscaling of meteorological data^{40,41}. In this context, conditional GANs are typically used, where both the generator and the discriminator are conditioned on low-resolution input fields in addition to the generator's noise input. We leverage the Wasserstein GAN (WGAN) with gradient penalty²¹, which is less prone to training instabilities such as mode collapse. The model architecture is based on prior work^{13,42}, and represents a commonly used model in downscaling.

Diffusion-based models

Diffusion models (DMs), in particular denoising score-matching approaches⁴³ are gaining traction in climate- and weather-modeling applications^{18,44,45}. In this work, we follow the diffusion-based downscaling framework introduced in ClimateDiffuse²², which combines several established components into an effective approach for conditional downscaling of climate fields. The model is a score-based model trained under the denoising score matching framework, in which a neural network is optimized to learn the score function. This score function is parameterized by a conditional U-Net, conditioned on low-resolution input fields, and the forward and reverse diffusion processes are defined via the stochastic differential equation formulation. Consistent with prior work²², we also incorporate several key design choices⁴⁶, such as the use of improved preconditioning and a higher-order integration scheme for the differential equation solver.

Training details

The GANs and DMs are trained for 60–168 hours (depending on training area size) with an effective batch size of 128 on 4 NVIDIA A100 GPUs. The ResNets are trained for 45–280 hours on 1 NVIDIA RTX8000. The years 2019 and 2020 are used as validation data for hyperparameter tuning and choosing the best checkpoint.

Evaluation

To evaluate spatial generalization ability of the models, we compare their performance across the different training scenarios using probabilistic and deterministic verification. Model performance is reported as absolute scores, relative improvements over interpolation baseline and relative performance compared to training directly on the target area.

Quantitative evaluation

We use the Continuous Ranked Probability Score (CRPS)⁴⁷ for probabilistic verification, a widely established metric that captures both, accuracy and calibration of the full predictive distribution. For grid-point-wise evaluation of extremes, we complement this with the percentile Mean Absolute Error (MAE). Inspired by prior works^{45,48}, we use this metric to directly evaluate the error in the value of the upper percentile.

We compute a point-wise CRPS using 8 samples at each of the six target locations. For a given forecast probability distribution F and the observed outcome y , the CRPS is defined as follows:

$$\text{CRPS}(F, y) = \int_{-\infty}^{\infty} [F(z) - 1(z \geq y)]^2 dz. \quad (1)$$

Here, $F(z)$ is the cumulative distribution function of the forecast distribution at point z and $1(\cdot)$ the indicator function. For a deterministic forecast, the CRPS reduces to the mean absolute error.

We compute the absolute error in the p th percentiles of the predicted and reference distributions as:

$$\text{MAE}_p = \frac{1}{D} \sum_{d=1}^D |x_{\text{pred}(\lfloor pN \rfloor, d)} - x_{\text{ref}(\lfloor pN \rfloor, d)}|, \quad (2)$$

where $\lfloor pN \rfloor$ gives the index corresponding to the p th percentile and $x_{(i),d}$ denotes the i -th order statistic (i.e., the i -th smallest value) of values for dimension d . We choose the 95th and 99th percentiles for computing the metric and average the values across all dimensions.

In-area training

The standard evaluation in deep learning-based downscaling consists of training and evaluating in the same area. To show that the geographical generalization is indeed a challenge we also report scores that use the target area labels. For this, we train the above mentioned models, ResNets, GANs and DMs on the respective target areas $E_1, \dots, E_6$ directly, while keeping the temporal train-test split: training on years 2001–2020 and testing on years 2021 and 2022.

Qualitative evaluation

In addition to CRPS, we provide a qualitative comparison of downscaled precipitation fields for one sample (see Fig. 9) as well as temporally averaged precipitation and CRPS values.

Geographic generalization through domain adaptation

Geographic generalization is a central challenge in statistical downscaling, resulting from differences in climatic conditions and their underlying processes across geographical areas. Such differences can lead to substantial performance drops when models trained in one region are applied to new regions with distinct climatic characteristics. While we include several predictors, such as orography or pressure, to capture local climatic conditions that shape precipitation through processes like orographic lifting, local precipitation is not determined by these factors alone. Precipitation dynamics, i.e. the processes that govern the intensity and variability of precipitation, also differ across regions. These dynamics are correlated with local climate conditions but are not fully determined by them. Instead, they are also influenced by large-scale circulation systems. Because of this interplay, regions with similar climate conditions may still experience very different precipitation regimes. These will not be captured by the predictors requiring additional strategies to address distribution shifts.

In our experiments, we observe large performance drops in target regions whose precipitation distributions differ markedly from those of the training regions (compare performance in Table 1 with Fig. 10). To address these distributional shifts, we explore two domain adaptation strategies. In the zero-shot setting, we apply quantile mapping to align precipitation distributions between training and target regions, and in the few-shot setting, we fine-tune pretrained models on a small subset of labeled target data.

Quantile mapping for geographical generalization in the zero-shot setting

While quantile mapping (QM) is traditionally used to correct systematic distributional biases in climate model simulations relative to observations (leveraging historical relationships between simulations and observations to adjust future simulations)⁴⁹, we apply it to address mismatches between the input distributions of the training and target regions. By aligning the input distribution of the target region more closely with that of the training data, the model may be better able to generalize and generate more reliable predictions in unseen regions. Specifically, we learn a mapping between the cumulative distribution functions (CDFs) of the precipitation input data of the training regions, $F_{\text{train},h}$, and that of the target regions, $F_{\text{target},h}$, over a historical period h . This results in the following transfer function:

$$\hat{x}_{\text{target},f}(t) = F_{\text{train},h}^{-1} (F_{\text{target},h} [x_{\text{target},f}(t)]) ,$$

which adapts the precipitation value $x_{\text{target},f}(t)$ from a target region within some future period f . For a non-negative variable such as precipitation we use the multiplicative variant of quantile mapping⁴⁹, where the values are lower bounded by zero. We perform quantile mapping to the low-resolution precipitation inputs using 1000 quantiles prior to data normalization during inference. Figure 11 illustrates the discrepancy between the CDFs

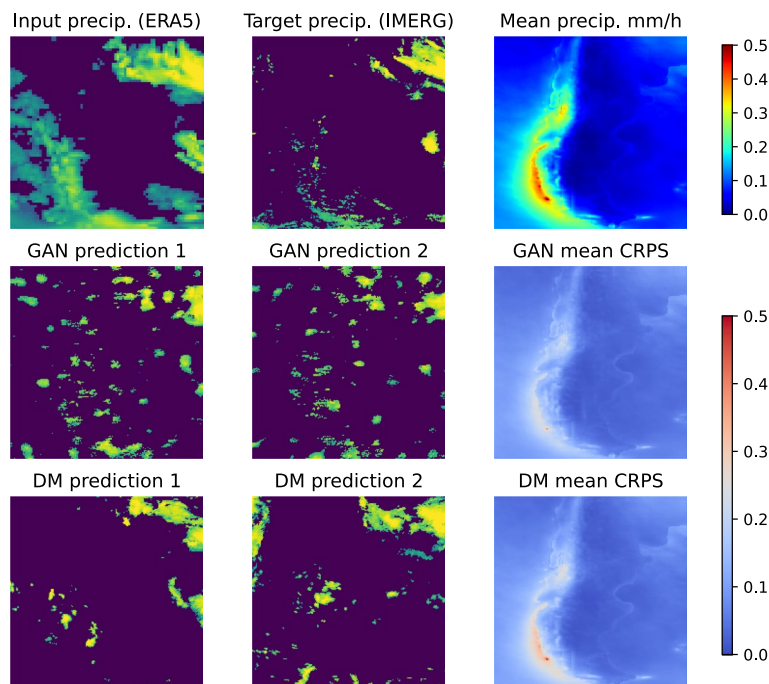


Fig. 9. Qualitative comparison of downscaled precipitation fields. This plot shows a sample, one time step from the evaluation set in the Cape Horn area in the first two columns and aggregated features in the last column. The random sample includes the matching input (ERA5) and target (IMERG) precipitation values in logarithmic scale and two samples each from the two generative approaches, GANs and DMs. The right column shows on top, the mean precipitation over the whole testing period (2021–2022) as well as the spatial distribution of CRPS values for GAN and DM.

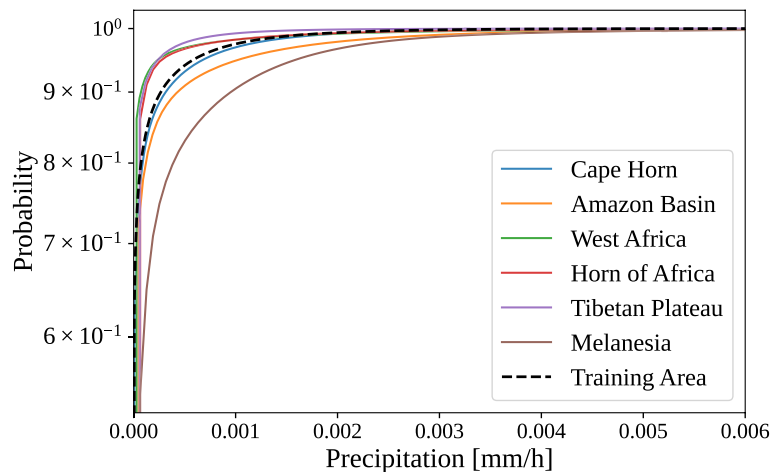


Fig. 10. Cumulative distribution functions of precipitation in training and target regions. The plot shows the empirical CDF of low-resolution precipitation from the training area (dashed curve) alongside the corresponding CDFs from all target regions. The plot shows clear differences in the precipitation input distributions across regions. For example, the CDF for Tibetan Plateau deviates strongly from that of the training region, showing mostly zero precipitation with little variance.

of the training and target regions, which is pronounced for all target regions except Cape Horn. After correction, the target region CDFs align more closely with the training CDF. Consistent with this reduction in distributional mismatch, we observe in Table 2 that all models evaluated with quantile-based aligned inputs achieves improved performance in most target regions (except Cape Horn) compared to using unaligned inputs.

Fine-tuning model to target regions in the few-shot setting

We, furthermore, consider the case where a small number of high-resolution samples for the target region do exist. In this few-shot setting, a model that is pre-trained on a large-scale dataset can be fine-tuned to a small-scale dataset as a form of transfer learning. To simulate this, we perform selected experiments using 2,000 randomly chosen data points from the target region. We focus on model-scenario-target region combinations that show the largest performance drops relative to in-distribution training (see Fig. 7). Following an approach used in prior work⁵⁰, we fine-tune the pre-trained model for a single epoch using a learning rate reduced to one-fifth of the original (0.0001). As shown in Table 3, even a small number of high-resolution samples from the target region can substantially improve model performance.

Domain adaptation resolves instabilities over the Tibetan Plateau

In Table 1 we observe an outlier value for the ResNet model trained on configuration A1 when evaluated over the Tibetan Plateau. While we consider this an isolated anomaly rather than a general trend, we investigate its potential cause. Since the model training itself remains stable, we hypothesize that the instability is a result of the large discrepancy in the precipitation distribution between training and target regions. As shown in Fig. 10, input precipitation for Tibetan Plateau is mostly zero with little variance, and normalization with training area statistics compresses these values into a narrow range around zero. Consequently, some samples may lie far outside the training distribution, causing the model (which is generally poor at extrapolating) to predict a few values that fall slightly outside the expected range. When the inverse transformation is applied via exponentiation, these deviations inflate into extreme outlier values (see Fig. 12 (right)). In Tables 2 and 3, we observe that both domain adaptation techniques—quantile mapping and fine-tuning—effectively resolve this effect by either aligning the precipitation inputs between training and target regions or adapting the pretrained model to the distributional shift.

Background on generalization and benchmarking in climate applications

Model generalization to new geographic regions is an active research area in applied machine learning, particularly in remote sensing and biodiversity modeling. In agricultural classification and segmentation, approaches such as task-informed meta-learning⁵¹, versions of model-agnostic meta-learning²⁶, and multi-source unsupervised domain adaptation²⁵ have shown promise in adapting to new regions with minimal data. In biodiversity monitoring, previous work integrates remote sensing and citizen science data to improve generalization in data-sparse regions like Kenya⁵², while spatial implicit neural representations have been leveraged for scalable global species range estimation using noisy, sparse data⁵³.

In climate science, spatial generalization in deep learning-based downscaling remains underexplored. A few recent studies have tested model transferability between subregions⁵⁴. For instance, some works examine generalization between different areas on the US West Coast^{8,35}. Others analyze performance across regions in the UK and the United States¹⁴ or evaluate generalization from the DACH region (Germany, Austria, and

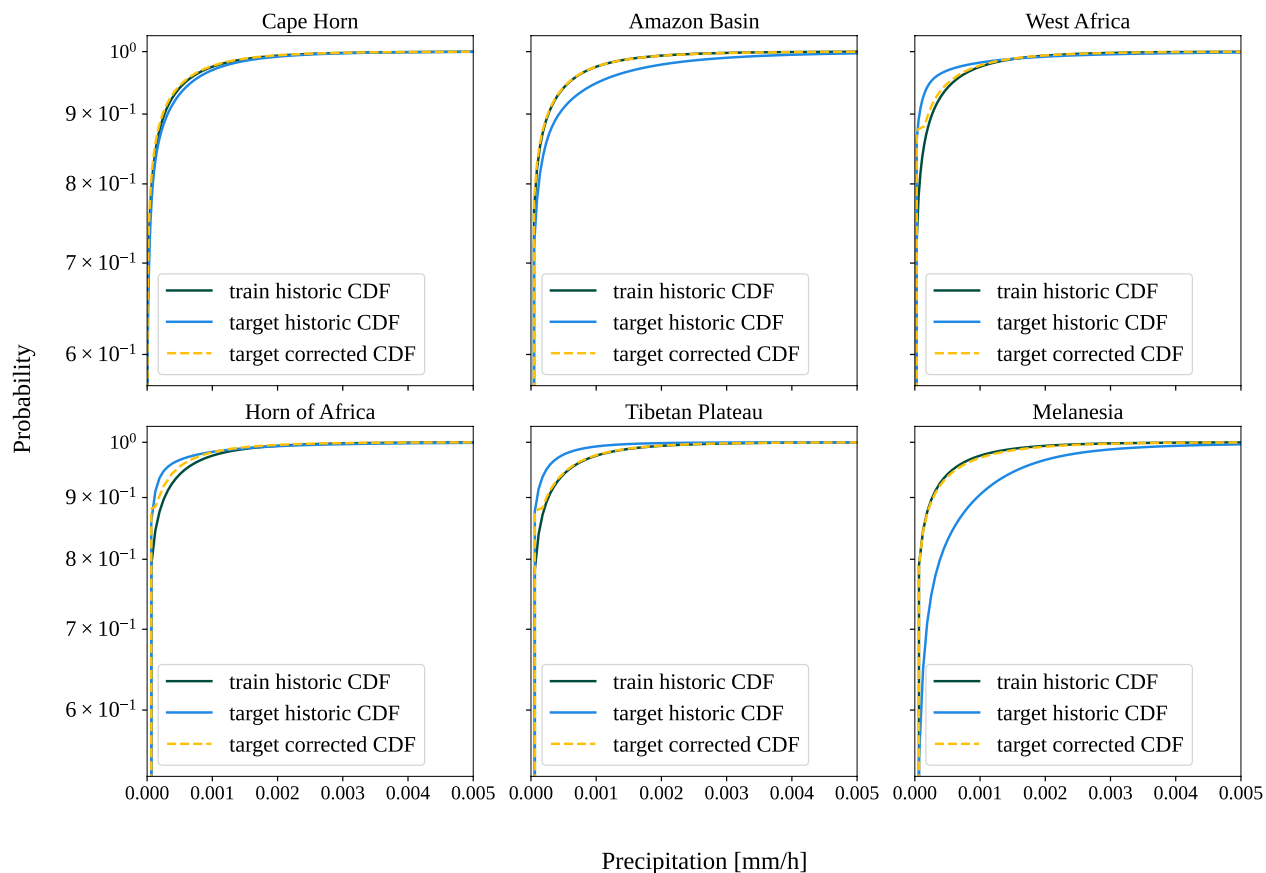


Fig. 11. Cumulative distribution functions (CDFs) of precipitation inputs for training and target regions. For each target region, a mapping is constructed between the historic CDF of the training region and the historic precipitation inputs of the target region. This mapping is subsequently used to correct the future input data from the target region, aligning them more closely with the training data distribution. Since precipitation is highly right-skewed with many zero and near-zero values, which causes most values to fall within the first bin of the CDF, we show the logarithm of the CDF to better visualize the discrepancies in distributions.

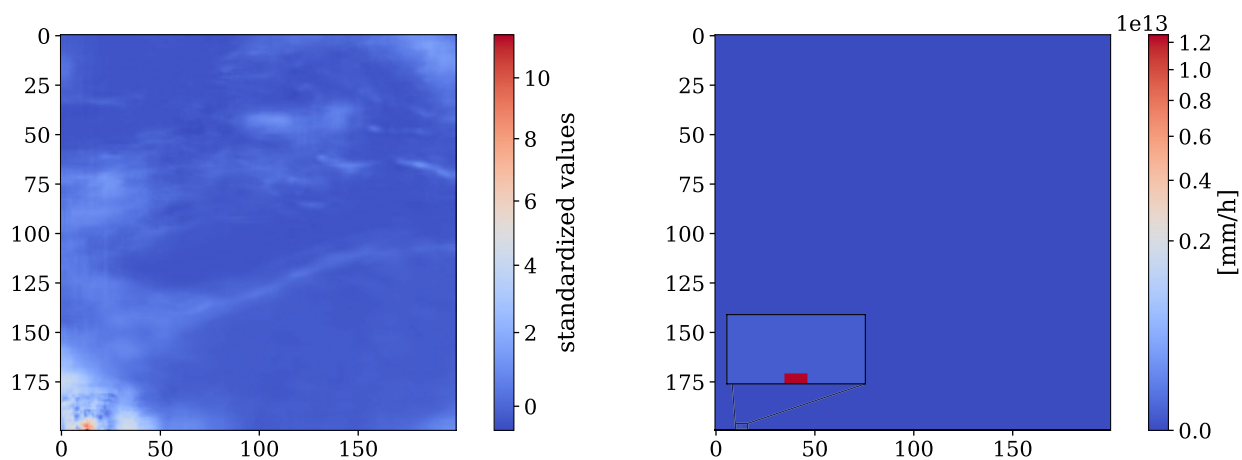


Fig. 12. Outlier ResNet prediction (single sample) over Tibetan Plateau. The abnormally large CRPS value is driven by a few single samples / pixel values that fall slightly outside the expected normalized range (left). Due to re-transformation via exponentiation such a pixel value gets largely amplified producing an extrem outlier pixel value (right).

Switzerland) to North America^{41,56}. While these efforts provide valuable insights, their geographic scope remains limited.

Another important development has been the creation of benchmark datasets aimed at standardizing and accelerating machine learning methods for Earth System Modeling. Notable examples include WeatherBench⁵⁷, ClimateBench⁵⁸, ClimateLearn⁵⁹ and ClimateSet⁶⁰ which have provided standardized datasets, tasks and evaluation frameworks. Several benchmarks specifically target precipitation forecasting and downscaling. RainBench introduces a global benchmark for precipitation forecasting based on IMERG data⁶¹. RainNet focuses on precipitation super-resolution, targeting a region on the US East Coast with single-variable input data in a deterministic setup⁶². The ClimateLearn benchmark supports evaluation of downscaling techniques that map low-resolution CMIP6 model outputs to high-resolution ERA5 data⁵⁹. These efforts have made climate modeling more accessible to the broader machine learning community. Despite this, no such existing benchmark has yet specifically addressed the challenge of generalizing across geographies. In this paper, we present RainShift, seeking to fill this gap by introducing a large-scale global benchmark dataset specifically designed to evaluate and improve the geographical generalization of deep learning-based downscaling.

Data availability

The RainShift benchmark dataset is hosted via Hugging Face at <https://huggingface.co/datasets/RainShift/rainshift>. The dataset consists of a zipped Zarr directory per region, resulting in overall 300GB of data. The download, data-loading, and training instructions can be found in our repository, that will be made available upon acceptance. All data used in this study are based on open-source datasets. The ERA5 reanalysis data are available from the Copernicus Climate Change Service (<https://cds.climate.copernicus.eu/datasets/reanalysis-era5-single-levels?tab=download>). The IMERG data are provided by NASA's Integrated Multi-satellite Retrievals for GPM and can be accessed at <https://gpm.nasa.gov/data/directory>.

Code availability

The code is available at <https://github.com/RolnickLab/downscaling-transferability>.

Received: 20 June 2025; Accepted: 29 December 2025

Published online: 27 January 2026

References

- Seneviratne, S. I. et al. Weather and climate extreme events in a changing climate. In *Climate Change 2021: The Physical Science Basis. Contribution of Working Group I to the Sixth Assessment Report of the Intergovernmental Panel on Climate Change* (ed. Seneviratne, S. I.) (Cambridge University Press, 2021).
- Ombadi, M., Risser, M. D., Rhoades, A. M. & Varadharajan, C. A warming-induced reduction in snow fraction amplifies rainfall extremes. *Nature* **619**, 305–310 (2023).
- Xiong, J. & Yang, Y. Climate change and hydrological extremes. *Curr. Clim. Change Rep.* **11**, 1 (2024).
- Maraun, D. et al. Precipitation downscaling under climate change: Recent developments to bridge the gap between dynamical models and the end user. *Rev. Geophys.* **48** (2010).
- Pendergrass, A. G., Knutti, R., Lehner, F., Deser, C. & Sanderson, B. M. Precipitation variability increases in a warmer climate. *Sci. Rep.* **7**, 17966 (2017).
- Fosser, G. et al. Convection-permitting climate models offer more certain extreme rainfall projections. *NPJ Clim. Atmos. Sci.* **7**, 51 (2024).
- Vandal, T. et al. DeepSD: Generating high resolution climate change projections through single image super-resolution. *Association for Computing Machinery* 1663–1672, <https://doi.org/10.1145/3097983.3098004> (2017).
- Sha, Y., Gagne, D. J., West, G. & Stull, R. Deep-learning-based gridded downscaling of surface meteorological variables in complex terrain. part i: Daily maximum and minimum 2-m temperature. *J. Appl. Meteorol. Climatol.* (2020).
- Höhllein, K., Kern, M., Hewson, T. & Westermann, R. A. comparative study of convolutional neural network models for wind field downscaling. *Meteorol. Appl.* **27**, e1961. <https://doi.org/10.1002/met.1961> (2020).
- Liu, Y., Ganguly, A. R. & Dy, J. Climate downscaling using ynet: A deep convolutional network with skip connections and fusion. In *Proceedings of the 26th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining, KDD '20*, 3145–3153, <https://doi.org/10.1145/3394486.3403366> (Association for Computing Machinery, New York, NY, USA, 2020).
- Rocha Rodrigues, E., Oliveira, I., Cunha, R. & Netto, M. Deepdownscale: A deep learning strategy for high-resolution weather forecast. In *2018 IEEE 14th International Conference on e-Science (e-Science)*, 415–422, <https://doi.org/10.1109/eScience.2018.00130> (2018).
- Goodfellow, I. J. et al. Generative adversarial networks. *Commun. ACM* **63**, 139–144 (2014).
- Harris, L., McRae, A. T. T., Chantry, M., Dueben, P. D. & Palmer, T. N. A generative deep learning approach to stochastic downscaling of precipitation forecasts. *J. Adv. Model. Earth Syst.* **14** (2022).
- Cooper, F. C., McRae, A. T. T., Chantry, M., Antonio, B. & Palmer, T. N. Further analysis of cgan: A system for generative deep learning post-processing of precipitation (2023). [arxiv:2309.15689](https://arxiv.org/abs/2309.15689).
- Sohl-Dickstein, J. N., Weiss, E. A., Maheswaranathan, N. & Ganguli, S. Deep unsupervised learning using nonequilibrium thermodynamics. *ArXiv abs/1503.03585* (2015).
- Mardani, M. et al. Generative residual diffusion modeling for km-scale atmospheric downscaling. *ArXiv abs/2309.15214* (2023).
- Wan, Z. Y. et al. Debias coarsely, sample conditionally: Statistical downscaling through optimal transport and probabilistic diffusion models. *ArXiv abs/2305.15618* (2023).
- Addison, H., Kendon, E., Ravuri, S., Aitchison, L. & Watson, P. A. Machine learning emulation of precipitation from km-scale regional climate simulations using a diffusion model (2024). [arxiv:2407.14158](https://arxiv.org/abs/2407.14158).
- Ling, F. et al. Diffusion model-based probabilistic downscaling for 180-year East Asian climate reconstruction. *npj Clim. Atmos. Sci.* **7**, 131. <https://doi.org/10.1038/s41612-024-00679-1> (2024).
- Tropical Globe. Tropical globe radar database. https://tropicalglobe.com/radar_database/ (2025). Accessed: 2025-01-26.
- Arjovsky, M., Chintala, S. & Bottou, L. Wasserstein GAN. *ArXiv abs/1701.07875* (2017).
- Watt, R. A. & Mansfield, L. A. Generative diffusion-based downscaling for climate (2024). [arxiv:2404.17752](https://arxiv.org/abs/2404.17752).
- Prasad, A. et al. Evaluating the transferability potential of deep learning models for climate downscaling. *ICML Workshop Machine Learning for Earth System Modeling* (2024).

24. Klemmer, K., Rolf, E., Robinson, C., Mackey, L. & Rußwurm, M. Satclip: Global, general-purpose location embeddings with satellite imagery. *ArXiv abs/2311.17179* (2023).
25. Wang, Y. et al. Exploring the potential of multi-source unsupervised domain adaptation in crop mapping using sentinel-2 images. *GI Sci. Remote Sens.* **59**, 2247–2265. <https://doi.org/10.1080/15481603.2022.2156123> (2022).
26. Rußwurm, M., Wang, S., Körner, M. & Lobell, D. Meta-learning for few-shot land cover classification. *Preprint arXiv 2004.13390* (2020).
27. Harder, P. et al. Hard-constrained deep learning for climate downscaling. *J. Mach. Learn. Res.* **24**, 1–40 (2023).
28. Hersbach, H. et al. The era5 global reanalysis. *Q. J. R. Meteorol. Soc.* **146**, 1999–2049. <https://doi.org/10.1002/qj.3803> (2020).
29. Hewson, T. & Pilloso, F. M. A low-cost post-processing technique improves weather forecasts around the world. *Commun. Earth Environ.* **2** (2020).
30. Huffman, G. et al. Integrated Multi-satellitE Retrievals for GPM (IMERG), version 4.4. NASA's Precipitation Processing Center (2014). Accessed: 31 March, 2015.
31. Huffman, G. J. et al. IMERG V07 Release Notes. <https://gpm.nasa.gov/resources/documents/imerg-v07-release-notes> (2024). Accessed: 2025-01-26.
32. Miles, A. et al. zarr-developers/zarr-python: v2.4.0, <https://doi.org/10.5281/zenodo.3773450> (2020).
33. Seyyedi, H., Anagnostou, E. N., Beighley, E. & McCollum, J. Hydrologic evaluation of satellite and reanalysis precipitation datasets over a mid-latitude basin. *Atmos. Res.* **164**, 37–48 (2015).
34. Xin, Y. et al. Evaluation of imerg and era5 precipitation products over the mongolian plateau. *Sci. Rep.* **12**, 21776 (2022).
35. Kumar, B. et al. Deep-learning based down-scaling of summer monsoon rainfall data over indian region. *Theor. Appl. Climatol.* (2021).
36. Hohllein, K., Kern, M., Hewson, T. & Westermann, R. A comparative study of convolutional neural network models for wind field downscaling. *Meteorol. Appl.* **27**, e1961 (2020).
37. Liu, Y., Ganguly, A. R. & Dy, J. Climate downscaling using ynet: A deep convolutional network with skip connections and fusion. In *Proc. 26th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, 3145–3153 (2020).
38. Sharma, S. C. M. & Mitra, A. Resdeepd: A residual super-resolution network for deep downscaling of daily precipitation over India. *Environ. Data Sci.* **1**, e19 (2022).
39. Rampal, N. et al. Enhancing regional climate downscaling through advances in machine learning. *Artif. Intell. Earth Syst.* **3**, 230066 (2024).
40. Rampal, N. et al. Enhancing regional climate downscaling through advances in machine learning. *Artif. Intell. Earth Syst.* <https://doi.org/10.1175/AIES-D-23-0066.1> (2024).
41. Glawion, L., Polz, J., Kunstmann, H., Fersch, B. & Chwala, C. Global spatio-temporal era5 precipitation downscaling to km and sub-hourly scale using generative ai. *npj Clim. Atmos. Sci.* **8**, 219 (2025).
42. Leinonen, J., Nerini, D. & Berne, A. Stochastic super-resolution for downscaling time-evolving atmospheric fields with a generative adversarial network. *IEEE Trans. Geosci. Remote Sens.* **59**, 7211–7223 (2020).
43. Sohl-Dickstein, J., Weiss, E., Maheswaranathan, N. & Ganguli, S. Deep unsupervised learning using nonequilibrium thermodynamics. In Bach, F. & Blei, D. (eds.) *Proceedings of the 32nd International Conference on Machine Learning*, vol. 37 of *Proceedings of Machine Learning Research*, 2256–2265 (PMLR, Lille, France, 2015).
44. Price, I. et al. Probabilistic weather forecasting with machine learning. *Nature* **637**, 84–90. <https://doi.org/10.1038/s41586-024-08252-9> (2025).
45. Hess, P., Aich, M., Pan, B. & Boers, N. Fast, scale-adaptive and uncertainty-aware downscaling of earth system model fields with generative machine learning. *Nature Machine Intelligence* 1–11 (2025).
46. Karras, T., Aittala, M., Aila, T. & Laine, S. *Elucidating the design space of diffusion-based generative models* (In Proc, NeurIPS, 2022).
47. Broecker, J. & Smith, L. A. *Increasing the Reliability of Reliability Diagrams*, vol. 22 (Weather and Forecasting, 2007).
48. Wan, Z. Y. et al. Regional climate risk assessment from climate models using probabilistic machine learning (2025). [arxiv:2412.08079](https://arxiv.org/abs/2412.08079).
49. Cannon, A. J., Sobie, S. R. & Murdock, T. Q. Bias correction of gcm precipitation by quantile mapping: how well do methods preserve changes in quantiles and extremes?. *J. Clim.* **28**, 6938–6959 (2015).
50. Zhang, Y. et al. Skilful nowcasting of extreme precipitation with nowcastnet. *Nature* **619**, 526–532 (2023).
51. Tseng, G., Kerner, H. & Rolnick, D. Timl: Task-informed meta-learning for agriculture (2022). [arxiv:2202.02124](https://arxiv.org/abs/2202.02124).
52. Teng, M. et al. Satbird: a dataset for bird species distribution modeling using remote sensing and citizen science data. In *Advances in Neural Information Processing Systems* Vol. 36 (eds Oh, A. et al.) 75925–75950 (Curran Associates, Inc., 2023).
53. Cole, E. et al. Spatial Implicit Neural Representations for Global-Scale Species Mapping. In *ICML* (2023).
54. Zhu, H. & Zhou, Q. Advancing satellite-derived precipitation downscaling in data-sparse area through deep transfer learning. *IEEE Trans. Geosci. Remote Sens.* **62**, 1–13. <https://doi.org/10.1109/TGRS.2024.3367332> (2024).
55. Sha, Y. II. & D. J. G., West, G. & Stull, R., Deep-learning-based gridded downscaling of surface meteorological variables in complex terrain. part ii: Daily precipitation. *Journal of Applied Meteorology and Climatology* **59**, 2075–2092. <https://doi.org/10.1175/JAMC-D-20-0058.1> (2020).
56. Prasad, A. et al. Evaluating the transferability potential of deep learning models for climate downscaling (2024). [arxiv:2407.12517](https://arxiv.org/abs/2407.12517).
57. Rasp, S. et al. Weatherbench 2: A benchmark for the next generation of data-driven global weather models (2023). [arxiv:2308.15560](https://arxiv.org/abs/2308.15560).
58. Watson-Parris, D. et al. Climatebench v1.0: A benchmark for data-driven climate projections. *J. Adv. Model. Earth Syst.* **14** (2022).
59. Nguyen, T., Jewik, J., Bansal, H., Sharma, P. & Grover, A. Climatelearn: Benchmarking machine learning for weather and climate modeling. *ArXiv abs/2307.01909* (2023).
60. Kaltenborn, J. et al. Climateset: A large-scale climate model dataset for machine learning. *ArXiv abs/2311.03721* (2023).
61. Towards data-driven global precipitation forecasting from satellite imagery. Schroeder de Witt, C. et al. Rainbench. *Proc. AAAI Conference on Artificial Intelligence* **35**, 14902–14910 (2021).
62. Chen, X. et al. Rainnet: A large-scale imagery dataset and benchmark for spatial precipitation downscaling. *Neural Inf. Process. Syst.* (2020).

Author contributions

P.H., D.R., C.L. and M.C. conceived the research. P.H., L.S., N.L., C.L., M.C. and D.R. designed the experiments. P.H. and F.P. implemented the code base and performed model training. P.H. and L.S. conducted the experiments and interpreted the results. P.H. and L.S. wrote the initial draft of the manuscript. P.H., A.H., and D.R. revised the manuscript.

Funding

Open Access funding enabled and organized by Projekt DEAL. This work was supported in part by the Canada CIFAR AI Chairs program and the Schmidt Sciences AI2050 Early-Career Fellows program. We also acknowledge computational support from Mila – Quebec AI Institute and the Digital Research Alliance of Canada, including in-kind support from Nvidia Corporation. The authors thank the International Max Planck Research School for Intelligent Systems (IMPRS-IS) for supporting Luca Schmidt.

Declarations

Competing interests

The authors declare no competing interests.

Additional information

Supplementary Information The online version contains supplementary material available at <https://doi.org/10.1038/s41598-025-34557-4>.

Correspondence and requests for materials should be addressed to L.S.

Reprints and permissions information is available at www.nature.com/reprints.

Publisher's note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

© The Author(s) 2026