

Machine-learned climate model corrections from a global storm-resolving model: Performance across the annual cycle

Anna Kwa¹, Spencer Koncius Clark², Brian Henn³, Noah D Brenowitz¹, Jeremy McGibbon¹, Oliver Watt-Meyer¹, W. Andre Perkins¹, Lucas Harris⁴, and Christopher S. Bretherton¹

¹Allen Institute for Artificial Intelligence

²Allen Institute for Artificial Intelligence / NOAA-GFDL

³Allen Institute for Artificial Intelligence (AI2)

⁴GFDL

September 14, 2022

Abstract

One approach to improving the accuracy of a coarse-grid global climate model is to add machine-learned state-dependent corrections to the prognosed model tendencies, such that the climate model evolves more like a reference fine-grid global storm-resolving model (GSRM). Our past work demonstrating this approach was trained with short (40-day) simulations of GFDL’s X-SHIELD GSRM with 3 km global horizontal grid spacing. Here, we extend this approach to span the full annual cycle by training and testing our machine learning (ML) using a new year-long GSRM simulation. Our corrective ML models are trained by learning the state-dependent tendencies of temperature and humidity and surface radiative fluxes needed to nudge a closely related 200~km grid coarse model, FV3GFS, to the GSRM evolution. Coarse-grid simulations adding these learned ML corrections run stably for multiple years. Compared to a no-ML baseline, the time-mean spatial pattern errors with respect to the fine-grid target are reduced by 6-25% for land surface temperature and 9-25% for land surface precipitation. The ML-corrected simulations develop other biases in climate and circulation that differ from, but have comparable amplitude to, the no-ML baseline simulation.

Machine-learned climate model corrections from a global storm-resolving model: Performance across the annual cycle

Anna Kwa¹, Spencer K. Clark^{1,2}, Brian Henn¹, Noah D. Brenowitz¹,
Jeremy McGibbon¹, Oliver Watt-Meyer¹, W. Andre Perkins¹,
Lucas Harris², and Christopher S. Bretherton¹

¹Allen Institute for Artificial Intelligence, Seattle, WA

²Geophysical Fluid Dynamics Laboratory, NOAA, Princeton, NJ

Key Points:

- Machine-learned (ML) corrective tendencies and radiative fluxes for a coarse-grid global model are trained using a year-long global storm-resolving simulation.
- The ML corrections reduce systematic errors of land surface temperature and precipitation in multiyear coarse-model simulations.
- The ML corrections nevertheless induce systematic biases in the simulated Hadley circulation.

Corresponding author: Anna Kwa, annak@allenai.org

Abstract

One approach to improving the accuracy of a coarse-grid global climate model is to add machine-learned state-dependent corrections to the prognosed model tendencies, such that the climate model evolves more like a reference fine-grid global storm-resolving model (GSRM). Our past work demonstrating this approach was trained with short (40-day) simulations of GFDL's X-SHIELD GSRM with 3 km global horizontal grid spacing. Here, we extend this approach to span the full annual cycle by training and testing our machine learning (ML) using a new year-long GSRM simulation. Our corrective ML models are trained by learning the state-dependent tendencies of temperature and humidity and surface radiative fluxes needed to nudge a closely related 200 km grid coarse model, FV3GFS, to the GSRM evolution. Coarse-grid simulations adding these learned ML corrections run stably for multiple years. Compared to a no-ML baseline, the time-mean spatial pattern errors with respect to the fine-grid target are reduced by 6-25% for land surface temperature and 9-25% for land surface precipitation. The ML-corrected simulations develop other biases in climate and circulation that differ from, but have comparable amplitude to, the no-ML baseline simulation.

Plain Language Summary

A recent vein of research uses machine learning to try and improve the predictive accuracy of climate models. Fine-resolution global storm-resolving simulations make more accurate rainfall and temperature predictions than climate models, but computers take too long to finish many-year simulations. We use data from a year-long high quality reference simulation to train machine learning models, which are then used in a lower quality but faster climate model. The machine learning applies corrections continually during the faster simulation to make it act more like the slow, high quality model. This improves the faster model's predictions for rainfall and temperature over land but also has unintended side effects of drying out the atmosphere and changing its circulation and tropical rainfall patterns.

1 Introduction

Due to computational constraints, running global climate models (GCMs) for more than a few years requires a spatial grid ($\gtrsim 50$ km) too coarse to resolve two key atmospheric physical processes: cumulonimbus convection and airflow over orography, coast-

lines and other land-surface heterogeneities. The subgrid variability of these processes are approximately represented in GCMs via expert-designed physical parameterizations. Subjective choices made within these parameterizations contribute significantly to uncertainty in GCM predictions of precipitation, cloud cover, etc (Woelfle et al., 2018; Zhao, 2014; Chen et al., 2007).

Global storm-resolving models (GSRMs), defined following Stevens et al. (2019) as global models with horizontal grid spacings < 5 km and at least 50 vertical levels, are able to better resolve these processes, but are currently too computationally demanding to be run for periods of more than a year or two. One way to improve the accuracy of GCMs while still maintaining the computational speedup from running at lower resolution is to train a machine learning (ML) model whose outputs can be applied to update the GCM state at each timestep, with the goal of making the GCM state evolve more like the coarsened state of an equivalent fine-grid GSRM model. Brenowitz & Bretherton (2019) and Yuval & O’Gorman (2020) trained ML models to predict the coarsened fine-grid apparent sources (Yanai et al., 1973) and apply these within coarse-grid aquaplanet simulations. Bretherton et al. (2022) (hereafter B22) trained ML models to predict corrections to the GCM subgrid parameterizations, which were then predicted and applied at runtime in a coarse-grid simulation with realistic topography.

Prior work has used the nudging framework as a means of generating training data and applying corrective machine-learned tendencies within a free-running coarse-grid simulation (Watt-Meyer et al., 2021; Bretherton et al., 2022; Clark et al., 2022). The nudging target varies across these works depending on the goal: nudging to observational reanalysis allows for effective bias correction without costly fine-grid simulations (Watt-Meyer et al., 2021), nudging to a higher-resolution (~ 25 km grid) model with parameterized convection allows for longer nudged runs to explore model performance in a range of climates (Clark et al., 2022), and nudging to a fine-grid (~ 3 km grid) GSRM allows the ML to benefit from a target dataset which directly resolves convection (B22).

B22 found that corrective ML models trained with the nudge-to-fine approach improved the weather forecast skill and reduced errors in time-mean land precipitation and surface temperature and the diurnal cycle of land precipitation in coarse-grid simulations of the 40-day period of their GSRM reference data. In this paper, we extend B22’s study by using a year-long GSRM reference simulation for training the corrective ML and eval-

uating simulations with the resulting ML-corrected coarse model on seasonal and multiyear timescales. In Section 2 we describe our coarse-grid and fine-grid models as well as the nudged coarse-grid simulation used as our training dataset. In Section 3 we describe the training procedure for our corrective ML models. Section 4 presents the offline (or single timestep) skill of the trained ML models. Section 5 presents online results and discusses the impact of the corrective ML when coupled to a free-running year-long coarse-grid FV3GFS simulation. Following B22, we use land precipitation and surface temperature as the primary metrics to gauge improvement over a free-running FV3GFS no-ML coarse simulation. We also discuss circulation changes in the ML-corrected simulations.

2 Simulations

2.1 Coarse-grid model

As in B22, the coarse-grid model is FV3GFS (Zhou et al., 2019) run at C48 (~ 200 km) resolution with the same 79 vertical model levels as the fine-grid model. Coarse simulations are carried out using a python-wrapped version of FV3GFS, which allows for easy customization and setup of nudging and integration of ML models (McGibbon et al., 2021). The model (and physics) timestep is 15 minutes, with 6 dynamics substeps per physics timestep. Time-varying sea surface temperatures (SSTs) and sea ice fraction are prescribed to be identical to those used in the coarsened fine-grid reference.

The following improvements were made to the coarse-model configuration in B22:

1. Fast saturation adjustment of humidity is enabled during each dynamics substep to better simulate fast-evolving cloud formation and dissipation. This significantly improves the baseline simulation’s surface radiation and precipitation climatology with respect to the fine-grid reference. Over a 40 day free-running coarse-grid simulation, the time-mean global-averaged surface downward shortwave and longwave radiative fluxes and precipitation root mean squared errors (RMSEs) are improved by 13, 12, and 30%, respectively.
2. The background vertical diffusion coefficient for heat and moisture is set to $2 \text{ m}^2/\text{s}$ to match the X-SHIELD value of this parameter over land. Without this, the default FV3GFS value of $1 \text{ m}^2/\text{s}$ led to crashes in the initial timestep.

We perform a free-running year-long simulation initialized from the coarsened fine-grid reference state on Jan. 19, 2020. This simulation has no ML corrections and is referred to hereafter as the “baseline” run.

2.2 Reference simulation

Our reference GSRM simulation is similar but not identical to the configuration used by B22, which was run for a much shorter 40-day period. In both this work and B22, the GSRM simulation is made with the X-SHiELD model, a modified version of FV3GFS with a C3072 cubed-sphere grid (~ 3 km spacing) and 79 vertical levels, run on NOAA’s Gaea computing system by collaborators at the Geophysical Fluid Dynamics Laboratory. The FV3GFS deep cumulus parameterization is disabled, though the shallow cumulus convection scheme is left active. The convective gravity wave drag scheme is disabled.

Cheng et al. (2022) describes the exact X-SHiELD configuration used here and some general features of the simulation. The free-running simulation is initialized on 2019-10-20. The first three months are excluded as spin-up, resulting in a year-long reference dataset spanning the remaining time from 2020-01-19 through 2021-01-17. The following are notable configuration differences in the reference fine-grid model used here with respect to the configuration in B22:

- There is no meteorological nudging of atmospheric fields to analysis.
- It uses the newer, inline version of the GFDL microphysics (Zhou et al., 2022).
- It uses a mixed-layer ocean between 45°S - 45°N , nudged with a 15-day time scale to ECMWF SSTs from analysis. SSTs outside of those latitudes are prescribed to ECMWF SSTs.
- Orographic gravity wave drag and mountain blocking schemes are enabled.

For use as the reference target in nudged training simulations, the X-SHiELD data is averaged from its native C3072 resolution down to the same C48 resolution as the coarse model, using the pressure-level coarsening procedure described in B22. All three-dimensional fields needed to restart the model, as well as numerous two-dimensional diagnostic fields, are coarsened inline and saved every 3 hours (B22 saved this output every 15 minutes, which is too expensive for our nine-fold longer training simulation).

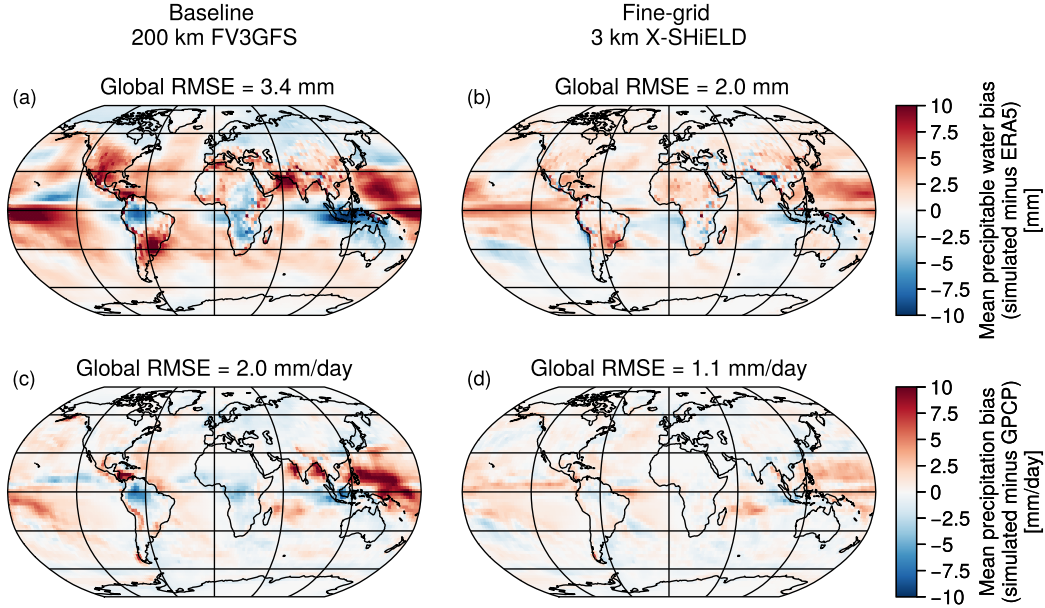


Figure 1: Top: Time-mean bias map of precipitable water over one simulated year in a coarse-grid FV3GFS run at C48 (~ 200 km) resolution (a) and storm-resolving X-SHIELD run at C3072 (~ 3 km) resolution (b). Bottom: Time-mean bias map of precipitation over one year in the coarse-grid FV3GFS (c) and fine-grid X-SHIELD (d) runs. The root mean squared error for each time-mean quantity is given in each subtitle.

Figure 1 compares coarse-grid FV3GFS and coarsened fine-grid X-SHIELD biases in 2020 annual-mean precipitable water with respect to ERA5 reanalysis (Hersbach et al., 2019) and precipitation biases with respect to GPCP observations (Adler et al., 2018). Both fields have a spatial RMSE with respect to the observations that is 40–45% lower in the 3 km X-SHIELD simulation than the coarser 200 km FV3GFS simulation. This supports using X-SHIELD as a reference target for improving the accuracy of our coarse-grid simulation.

2.3 Nudged training simulation

We follow the nudging framework of B22. We initialize a one-year coarse-grid FV3GFS simulation with the configuration described in Section 2.1 using the coarsened X-SHIELD state on 2020-01-19. Temperature, humidity, model layer pressure thickness, and horizontal wind fields are nudged at each model time step towards the fine-grid state with

a 3-hour nudging timescale and the interval-averaged nudging tendencies at all coarse grid points are saved every three hours. The reference state is interpolated for timesteps that lie between the 3-hour intervals of the fine-grid data. The nudging tendencies are calculated as follows in Equation 1:

$$\Delta Q_a = -\frac{a^n - \bar{a}}{\tau} \quad (1)$$

where a^n is a prognostic field in the model, $\bar{a}$ is the coarsened value of that field in the reference fine-resolution data, and τ is a constant nudging timescale. As in B22, we use a nudging timescale τ of 3 hours. We interpret the nudging tendencies as corrections that would make the coarse model follow the evolution of the coarsened fine grid model and machine-learn these corrective tendencies as functions of the column state.

We use the following approach from B22 to correct for systematic biases in surface precipitation and downwelling radiative fluxes due to the coarse model physics producing less cloud and land precipitation than the fine-grid model. During the nudged training run we prescribe surface downwelling shortwave and longwave fluxes as well as precipitation from the coarsened fine-grid output to the land model. This prevents biases in land surface properties from feeding back into the atmosphere and affecting the temperature and humidity nudging tendencies. We train a ML model to predict downward longwave and shortwave surface fluxes for application in prognostic simulations. The radiative flux model training scheme is modified from B22 to improve its accuracy, as described below.

3 Machine-learned corrective models

3.1 Training and test data

The nudged run data is divided into interleaved blocks of two weeks of training data followed by one week of validation data. The first and last two timesteps of each two-week training data block are discarded to ensure that there are no consecutive timesteps in both the training and validation datasets.

1367 timesteps are selected out of the available training data. We make the assumption that the corrections predicted by the neural networks are column-local and only depend on the state within a single grid column. At C48 grid resolution, each timestep con-

tains 13824 columns. Since nearby columns are strongly correlated, inclusion of all columns in a timestep does not provide significant benefit over using a reasonable subsample of the data. Initial experiments using a training dataset of 200 timesteps showed that subsampling down to about 10% of the global set of columns in each timestep did not negatively impact the validation loss compared to using all columns. We use a subsample fraction of 15% of the total number of columns in each timestep in order to reduce memory usage and training time. Thus the training dataset consists of $\sim 2.8 \times 10^6$ samples. The test dataset used for offline evaluation is subsampled to 100 of the available test timesteps, with no column subsampling so as to allow for the creation of time-mean offline bias maps.

3.2 Neural network training

Following B22, we train two fully connected dense neural networks (NNs) to separately predict i) vertical profiles of air temperature and specific humidity tendencies and ii) column shortwave transmissivity and downward longwave surface flux. We train four NNs for each set of outputs using different random seeds and also construct an ensemble model using all of the randomly seeded NNs which outputs the median prediction of the ensemble members for each field.

Choices of width, depth and learning rate were guided by a hyperparameter sweep in a randomized grid search (Biewald, 2020). To speed the hyperparameter sweep, we used the hyperband algorithm (Li et al., 2018) for early stopping, where training was terminated after 10 epochs if validation loss was not improved relative to previously tested sets of hyperparameters. Validation loss was primarily affected by the choice of learning rate; width and depth had much less impact on model skill. We first chose the value of learning rate that performed best in the sweep, and then set the network width and depth using the set of sweep parameters that performed best near that value of learning rate.

3.2.1 *Air temperature and specific humidity tendencies NN*

The NN trained to predict corrective air temperature and specific humidity tendencies is a fully-connected dense network with 3 hidden layers with 419 neurons per layer. It is trained for 500 epochs with early stopping and a learning rate of 1.4×10^{-4} . Its input features are the cosine of solar zenith angle, surface geopotential, latitude, and the

vertical profiles of air temperature and specific humidity. The surface geopotential implicitly provides information about the surface type (land or sea/sea ice). Latitude is included as a new feature not used by B22, because it improves offline model skill at high latitudes, where there are extreme conditions but relatively few columns contributing to the loss function.

As in B22, inputs to the network are normalized via standard scaling each vertical level using the mean and standard deviation of the first 4×10^5 samples. Output profiles passed to the loss function are normalized by the standard deviation of each vertical level such that all vertical levels are equally weighted in the loss. We use the same mean absolute error loss and L2 regularization penalty of 10^{-4} as B22.

The following new ML configuration options helped ensure online stability:

- ML-predicted tendencies of heating and humidity are limited to magnitudes less than 0.002 K/s and 1.5×10^{-6} kg/kg/s, respectively. These limiters are applied as a layer within the dense NN such that the limited outputs are used during optimization. These ranges comfortably extend beyond the nudging tendency minima and maxima in the training data by a factor of 3, but prevent the NNs from making extreme predictions when undesirable feedback between the coarse-grid model and ML corrections lead to atmospheric input states outside the envelope of the training data.
- Following Clark et al. (2022), we exclude (‘clip’) the uppermost 25 vertical model levels ($\lesssim 150$ hPa) of specific humidity and air temperature state inputs from the feature set. Without doing this, the models’ Jacobian matrices showed that the output fields in the boundary layer were as sensitive to inputs from the uppermost 25 model levels as they were to input levels in their immediate locality (Brenowitz & Bretherton, 2019).
- The uppermost 3 vertical levels of temperature and humidity tendency outputs are excluded from the prediction. The ML model always applies a zero corrective tendency for these levels when used online. Differences in the sponge layer damping between FV3GFS and the X-SHiELD reference model lead to large nudging tendencies in these few levels with magnitudes similar to those in the boundary layer, but we do not consider these differences to be part of the coarse model physics that we wish to correct. The output clipping here is less extensive than approach

taken in Clark et al. (2022), where the output tendencies were tapered to zero at the top of the model in the uppermost 25 levels. We found that clipping just the uppermost three levels provided similar benefits in terms of online stability and performance.

3.2.2 *Surface radiative flux NN*

We train a NN to predict surface downward longwave radiative flux and column shortwave transmissivity to correct for the effect of systematic cloud biases on the land surface in the coarse run (Section 2.3). Its input features are the cosine of solar zenith angle, surface geopotential, latitude, and the vertical profiles of air temperature and specific humidity. Transmissivity is set to zero in the training data for nighttime columns with zero solar insolation. The predicted column transmissivity is multiplied by the top-of-atmosphere downward shortwave flux to infer the predicted downward shortwave flux at the surface. This approach, introduced by Clark et al. (2022), differs slightly from the radiative flux NN in B22, which directly predicted the downward and net surface shortwave radiative fluxes and did not include latitude as an input feature.

Like the tendency model, the surface radiative flux NN is a fully-connected dense network with 3 hidden layers of width 419, trained for 500 epochs with early stopping and mean absolute error loss. It uses a learning rate of 4.9×10^{-5} and L2 regularization penalty of 10^{-4} .

Longwave flux outputs are enforced to be positive or zero, and transmissivity outputs are limited to the range $[0, 1]$. As in the tendency NN, these limits are applied as a dense network layer such that the limited outputs used in the loss function.

3.3 Sensitivity to data sampling and NN configuration

We had to update our training methodology from B22 for the ML corrections to be stable and skillful over a yearlong simulation. Early NN versions using similar training dataset size and hyperparameters as in B22 developed large drifts over a year-long simulation, with some random seed variants crashing before a full year completed. A leading indicator of drift and instability was a steady warming of 200 hPa air temperature starting at the poles and growing to a ~ 20 K warm bias that extended into the mid-latitudes

within 90 days. This online bias was linked to a positive offline bias ML heating tendencies near the poles in the upper troposphere.

Several changes described in Section 3.2 improved our offline skill and subsequently reduced the online growth of stratospheric air temperature biases. In Table 1 we present the impact of those configuration choices on offline skill. The last column of Table 1 lists the polar regions’ mean offline bias in 150-400 hPa heating tendencies; though this metric is calculated offline we found it to be a useful proxy for online air temperature drifts in the first few months of the simulations.

Updates are described relative to a starting “base” configuration, which used the same dataset size as B22 (130 timesteps randomly chosen from the year of data, with no downsampling of columns) as well learning rate (0.002) and number of training gradient descent steps ($\sim 2.9 \times 10^7$). The input clipping, limiters on output magnitudes, and the clipping of the uppermost 3 vertical tendency levels (described in Section 3.2.1) are present in all the configurations tested below. Table 1 compares the following configuration choices, applied sequentially:

- + **lat** Latitude is included as an input feature.
- + **sampling** The number of timesteps in the training dataset is increased $>10\times$ and subsample down to 15% of the grid columns in each timestep.
- + **lower LR** The learning rate is lowered from the value used in B22 from 0.002 to 0.0014 and the number of gradient descent steps is increased $>30\times$.

Training with a lower learning rate over more gradient descent steps is the main contributor to the increased global offline R^2 of column-integrated tendencies.

4 Offline performance

Here we discuss the ML models’ “offline” skill in predicting their training data targets over a single time step. Offline skill is a necessary but not sufficient condition for successful application of the corrective ML in the coarse model. However, we do not strictly aim to maximize offline skill, as we find that regularization (which slightly lowers offline skill) is required for online stability.

As in Figure 4 of B22, the instantaneous nudging tendencies at any given time step are quite noisy, so the ML cannot be expected to have perfect skill. In Figure 2 we show

Configuration	$\langle \text{heating} \rangle_{\text{ML}}$	$\langle \text{moistening} \rangle_{\text{ML}}$	polar 150-400 hPa
	R^2	R^2	ML heating bias [K/s]
Base	0.17	0.15	1.1E-06
Base + lat	0.18	0.16	8.1E-07
Base + lat + sampling	0.18	0.15	7.0E-07
Base + lat + lower LR	0.28	0.22	-1.8E-07
Base + lat + sampling + lower LR	0.29	0.23	2.2E-07

Table 1: Sensitivity of offline column-integrated ML-predicted heating $\langle \text{heating} \rangle_{\text{ML}}$ R^2 , column-integrated ML-predicted moistening $\langle \text{moistening} \rangle_{\text{ML}}$ R^2 , and mean 150-400 hPa ML heating bias at the poles ($|\text{lat}| > 60^\circ$) to various changes in the data sampling and NN configuration.

the zonal and pressure-level mean coefficient of determination R^2 on the offline testing data for the ML-predicted corrective tendency fields. Both temperature and humidity tendency predictions are most skillful in the boundary layer and in the tropics, with zonal-mean R^2 values upwards of 0.8 in the tropical boundary layer and 0.5–0.8 in the tropical free troposphere and extratropical boundary layer. Model skill in the mid-to-upper troposphere degrades to 0.1–0.3 at higher latitudes. The global-mean vertical profiles of R^2 (not shown) are much improved relative to Figure 5 of B22, increasing up to a factor of two in the lower to mid-troposphere. This improvement in offline skill from B22 results from the updates to our training methodology and configuration described above.

As described in Sec. 3.2.2, the ML for downwelling surface radiation is slightly different from B22. Time-mean offline biases for the downward radiative fluxes are shown in Table 2. Their global average biases are small, regardless of neural-net random seed choice: $1.2 \pm 0.7 \text{ W/m}^2$ for downward shortwave and $-0.3 \pm 0.2 \text{ W/m}^2$ for downward longwave surface radiative flux. The RMSE of the time-averaged total downward flux prediction is significantly reduced from 11.6 W/m^2 in B22 to 3.4 W/m^2 here.

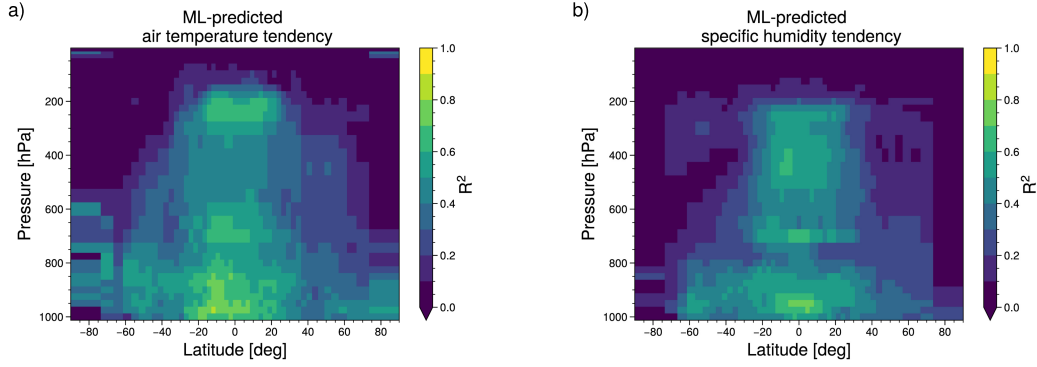


Figure 2: R^2 of offline predictions of air temperature and humidity tendencies, averaged over latitude and pressure. Predictions are generated using the NN ensemble; results for individual seeds are similar.

Surface radiative flux field	RMSE [W/m^2]	Bias [W/m^2]	R^2
Downward shortwave	3.8	1.2	0.99
Downward longwave	1.2	-0.3	0.99
Total downward	3.4	0.9	0.99

Table 2: Offline metrics for downward surface fluxes. Downward shortwave and longwave fluxes are predicted by the radiative flux NN; total downward flux is the sum of shortwave and longwave predictions.

5 Online performance

The true test of the corrective ML comes by applying it online in a prognostic simulation of the coarse-grid FV3GFS model and measuring results over seasonal and year-long timescales. Here, we present results from a suite of five such ML-corrected simulations. Four use independent random seeds to initialize the tendency and radiative flux NNs. A fifth uses the median of the corrective tendencies and surface fluxes over this four-member NN ensemble. All ML-corrected simulations ran stably for the full simulation length of 360 days.

For each one-year ML-corrected simulation we assess the global-mean biases and seasonal-mean spatial pattern errors of the precipitation and land surface temperature with respect to the X-SHIELD reference run. These are compared to the baseline coarse-grid FV3GFS simulation. All coarse-grid FV3GFS simulations use the same namelist configuration, initial conditions and prescribed sea ice and SSTs. Ideally, the biases and pattern errors of the ML-corrected simulations will be significantly smaller than those of the baseline simulation. We also examine the effects of including the ML correction on other measures of large scale circulation and climate drift in the coarse simulations.

5.1 Improvements in precipitation and surface temperature

Figure 3 displays the RMSE and bias of time-averaged surface precipitation over the simulation year. The ML corrections from the individual NNs and the NN ensemble all improve the coarse model’s time-mean precipitation skill over both land and ocean. As we only sample four randomly seeded models, we will cite the minimum and maximum relative improvement in each metric across the seeds instead of a standard deviation. ML-corrected models improve upon the baseline precipitation RMSE by 13–21% globally, 9 – 25% over land, and 13 – 20% over ocean. The magnitude of the precipitation bias is strongly reduced by 63–89% globally, 43–88% over land, and 82–98% over ocean.

The time-averaged precipitation error pattern is qualitatively similar across the various NNs. For conciseness we only show the ML-corrected simulation using the NN ensemble in Figure 4a. The baseline coarse run tends to have a large negative bias in precipitation over equatorial Africa and South America, and a large positive bias over the western Pacific warm pool. The ML-corrected model reduces the magnitude of these re-

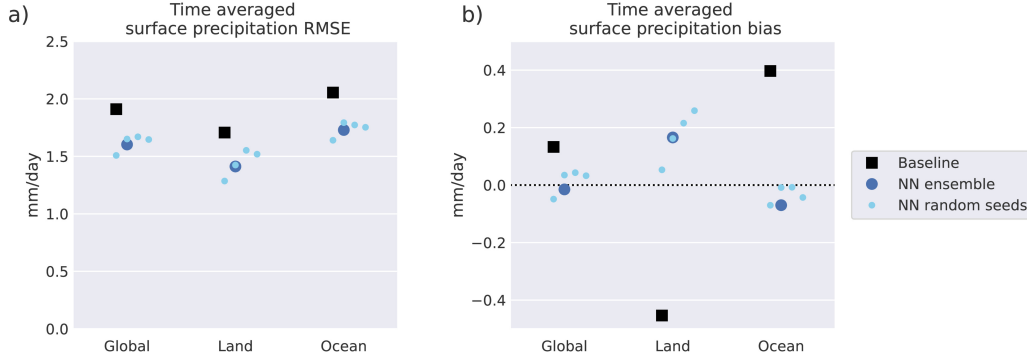


Figure 3: Time-mean RMSE and bias (with respect to the fine-grid reference) of precipitation in the baseline and ML-corrected coarse runs, shown for global, land, and sea domain averages.

gional biases. The baseline’s wet bias in the western Pacific is replaced with a dry bias of slightly smaller magnitude. The pattern error of precipitation in the ML runs in the tropics differs significantly from the baseline due to the changes in circulation brought about by the ML correction (see section 5.2).

Figure 4b shows the seasonally averaged errors in precipitation over land in the baseline and ML-corrected runs. The improvement in land precipitation errors is robust across all seasons in all four NNs as well as the NN ensemble.

The 10–20% relative improvement over the baseline in the RMSE of time-averaged surface precipitation is less than the $\sim 30\%$ reduction found by B22. We discovered that enabling fast saturation adjustment on each of the six dynamics substeps within the dynamical core reduced the baseline model precipitation RMSE from 3.7 mm/day over the 40 day run in B22 down to 1.9 mm/day in annual average here. Thus, the absolute errors in the baseline as well as the ML-corrected runs are notably improved in this work over B22.

Land surface temperature errors are also reduced by 6–26% in the ML-corrected runs. Figure 5a shows the annual-average land surface temperature pattern error in the baseline and NN ensemble runs. The ML-corrected runs reduce a warm bias across Africa and the western United States. Seasonal surface temperature RMSE is plotted in Figure 5b. While ML-corrected runs consistently improve land surface temperature from

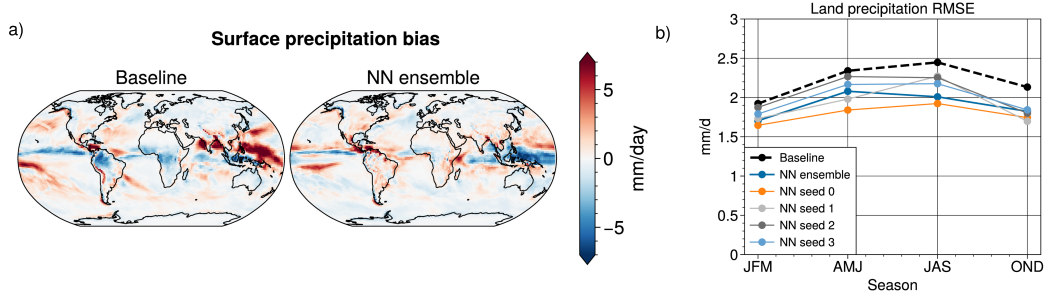


Figure 4: Left: Time-mean error pattern map of surface precipitation relative to the fine-grid model. The ML-corrected run uses the NN ensemble. Right: Seasonal RMSE of land surface precipitation for the coarse model baseline and ML-corrected simulations.

April through September, their behavior in boreal winter is less consistent across seeds, with NNs generated from two seeds performing comparably or worse than the baseline in these months. Those simulations amplify a systematic cold bias during boreal winter at high northern latitudes ($\gtrsim 50^\circ$ N) also present in the baseline, while ML-corrected runs with other seeds reduce this bias (not shown).

Both land surface temperature and precipitation show the largest seasonal skill improvements during April - September, with the best-performing model (seed 0) reducing both land temperature and precipitation seasonal RMSEs by over 20% in these seasons.

5.2 Other climate biases in ML-corrected simulations

B22 found that 40-day coarse-grid simulations with and without corrective nudging both developed mean-state biases in the latitudinal structure of upper-tropospheric (200 hPa) temperature which could be 5 K or more in parts of the extratropics (their Figure 13d). The full year of reference X-SHiELD simulation data allows us to test how these and other biases develop over longer timescales. We compare just the NN ensemble model to the baseline, as the biases discussed here are robust across all ML-corrected simulations.

Our ML-corrected simulations develop a warm bias in 200 hPa air temperature (Figure 6a). This bias is no more than ~ 5 K over most of the globe, but is substantially larger

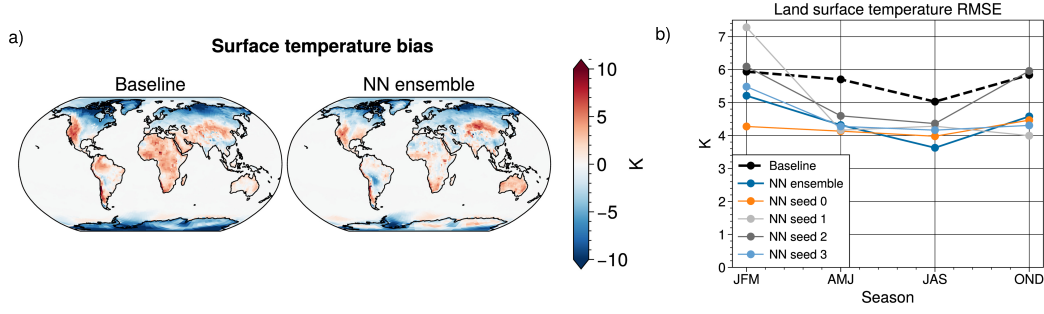


Figure 5: Left: Time-mean error pattern map of land surface temperature relative to the fine-grid model. The ML-corrected run uses the NN ensemble. Surface temperature is only shown over land and sea ice, as the sea surface temperature is prescribed from the fine-grid reference. Right: Seasonal RMSE of land surface precipitation for the coarse model baseline and ML-corrected simulations.

at high southern latitudes. In contrast to B22, all the NNs in this work drift similarly over the year-long run. The baseline run develops a smaller cold bias at most latitudes.

The ML-corrected runs also develop robust bias patterns in precipitable water (Figure 6b). A strong dry bias of up to -10 mm develops over the seasonally shifting GSRM-simulated tropical ocean ITCZs, as well as a northern summertime dry bias at high latitudes northwards of $\sim 50^\circ$. The Northern Hemisphere dry bias is also evident in the baseline run, albeit to a lesser magnitude. The baseline model does not share the dry ITCZ bias, but is too moist in most other parts of the subtropics.

A related bias of the ML-corrected runs is a weakened Hadley circulation that is also shifted southward during boreal spring and summer. Figures 6 c-d show the biases in precipitation and 500 hPa vertical wind in the baseline and ML-corrected runs with respect to the coarsened fine-grid reference. The baseline model has a northward shift of vertical wind and precipitation in the tropics relative to the fine-grid reference during Jun-Aug. In contrast, the ML-corrected run displays southward-displaced upward motion and precipitation during Jun-Aug.

We do not yet have a remedy for these ML-induced climate biases. Like B22, we tried including a NN that learned corrective tendencies for the zonal and meridional winds. B22 found that including ML-corrected wind tendencies led to large time-mean 200 hPa

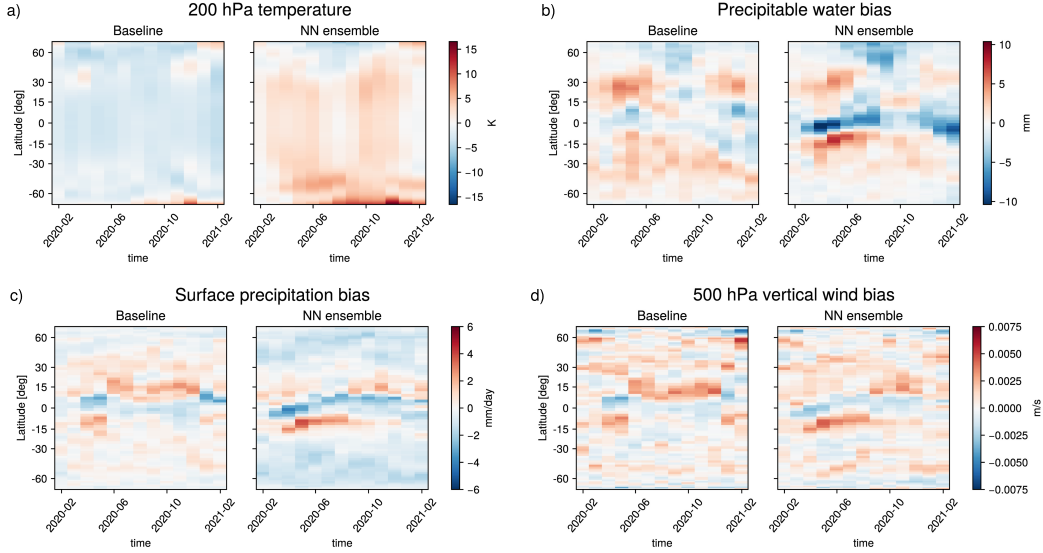


Figure 6: Zonal mean vs. time plots for biases (with respect to the fine-grid model) in 200 hPa temperature, precipitable water, precipitation, and 500 hPa vertical wind in the baseline and ML-corrected prognostic runs.

air temperature biases developing over their 40-day run. In the present study, we similarly found that including corrective wind tendencies caused errors in 200 hPa temperature that grew quickly within the first 7 days of the run and ultimately led to numerical instability. We are currently exploring the use of additional online ML models to regulate the application of ML corrections and prevent predictions on out-of-sample states from pushing the model farther outside of the training data envelope. This looks to be a promising direction for reducing climate biases by incorporating wind corrections without triggering instabilities.

5.3 Tropical variability

This section documents that the space-time variability of tropical precipitation is a weak point of our ML-corrected simulations, even though they improve aspects of the seasonal mean precipitation. We believe this is mainly caused by training using nudging tendencies. Because the coarse-grid convective parameterization includes a trigger that is highly sensitive to the column thermodynamic state, the precipitation predicted by physical parameterizations of the nudged coarse simulation is intermittent and poorly

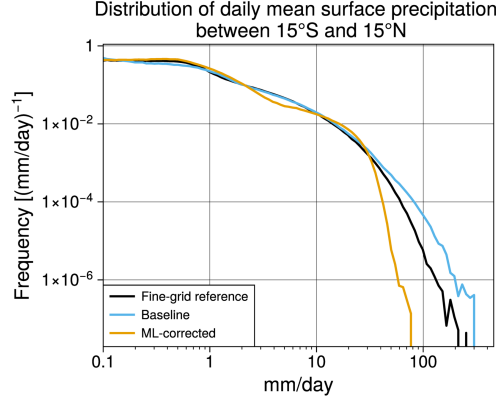


Figure 7: Distribution of tropical ($15^{\circ}\text{S} - 15^{\circ}\text{N}$) daily mean surface precipitation in the fine-grid, baseline, and ML-corrected simulations.

synchronized with the more smoothly varying precipitation in the reference fine-grid model. The temperature and humidity nudging tendencies mediate this disagreement by damping the temperature and humidity tendencies in the coarse-grid physics; this is learned by the ML correction. Hence, the ML correction tends to reduce the sensitivity of the physical parameterizations to column thermodynamic state, reducing extreme precipitation and tropical variability of precipitation.

Figure 7 shows the PDF of daily precipitation between $15^{\circ}\text{S} - 15^{\circ}\text{N}$. The baseline run has excessive extreme precipitation but the ML-corrected simulation has too little extreme precipitation.

Figure 8 plots tropical ($15^{\circ}\text{S} - 15^{\circ}\text{N}$ average) precipitation across longitude and time in the fine-grid and free-running coarse-grid simulations. Eastward-propagating Kelvin waves and hints of a Madden-Julian Oscillation stand out in the fine-grid reference but are largely absent in both coarse-grid simulations. The baseline run has overly-strong precipitation in the western Pacific ($\sim 150^{\circ} - 200^{\circ}\text{E}$) and too much power in westward-propagating waves; in contrast, the ML-corrected model largely damps out zonally propagating waves in both directions and underpredicts western Pacific precipitation.

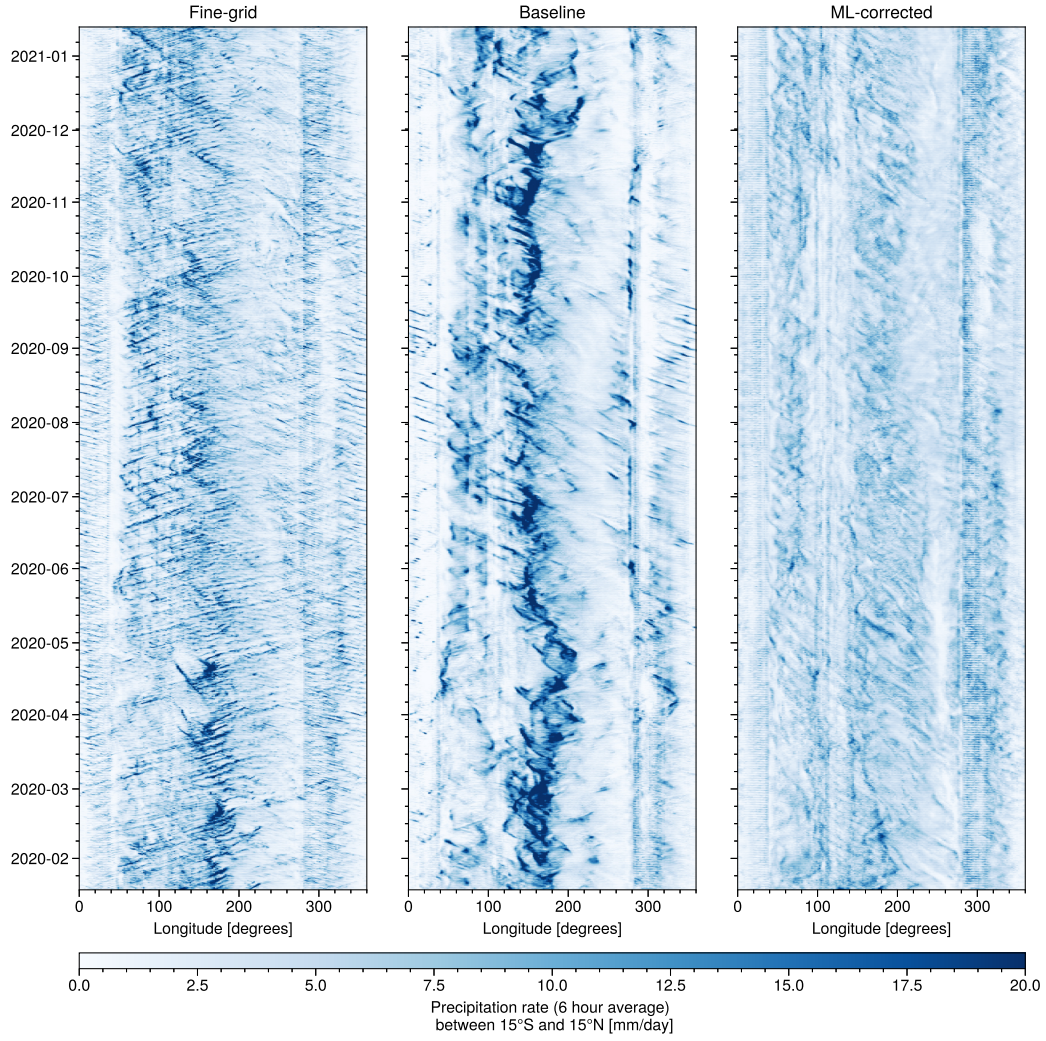


Figure 8: Longitude-time plots of 6-hourly averaged precipitation rate between $\pm 15^\circ$ latitude in the fine-grid, baseline, and ML-corrected simulations.

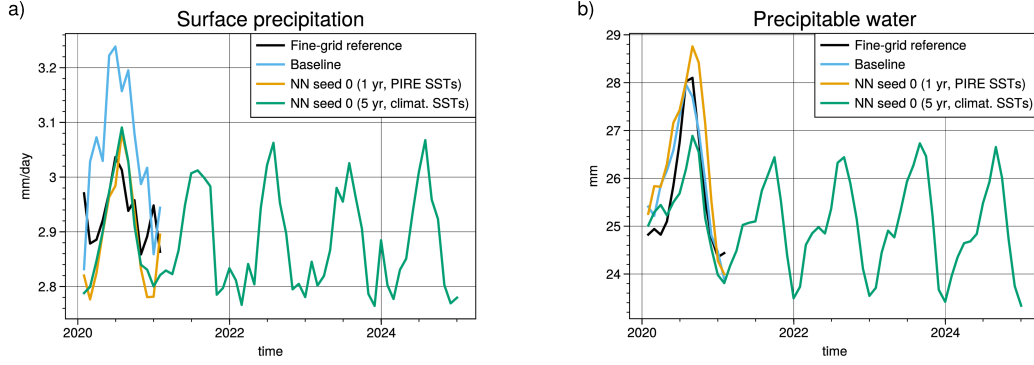


Figure 9: Globally-averaged time series of monthly-mean surface precipitation and precipitable water from the reference fine-grid, baseline coarse-grid, and ML-corrected coarse-grid simulations. Unlike the yearlong baseline and ML runs, the 5 year ML run uses climatological SSTs and is thus not limited to the time range of the reference run.

5.4 Stability over multiple years

Would the climate in ML-corrected runs continue to drift, with larger biases developing over multiple years of simulation, or are its biases largely repeatable in subsequent annual cycles, as in the no-ML baseline? This interannual stability under constant SST forcing is a prerequisite for use in climate-length simulations.

To address this question, we used the NN with the lowest surface temperature and precipitation errors (seed 0) for a 5-year coarse-grid run. All prognostic runs so far used the SSTs from the yearlong reference X-SHiELD simulation, which did not span the full length of the extended 5 year simulation and did not exactly repeat at the end of the annual cycle. To enable a smoothly-forced five-year simulation, we instead used a climatological annual cycle of SSTs.

Figure 9 shows time series of precipitable water and precipitation in this extended ML-corrected run as well as the yearlong fine-grid reference, baseline, and ML-corrected runs with SSTs prescribed from reference data. As desired, the multiyear ML-corrected run maintains consistent seasonal bias patterns across years that match up well in its first year to the fine-grid X-SHiELD reference simulation.

6 Conclusions

This study extends previous work done in B22 in which corrective ML models trained using fine-resolution data were applied within coarse-grid climate models. The novelty of this work lies in the training and evaluation of the corrective ML over the entire annual cycle. This advance was enabled by the use of the year-long X-SHiELD simulation reference dataset. The longer-term goal is to use GSRM simulations in multiple climates (Cheng et al., 2022) to train corrective ML that can be used in climate change simulations, following the template in Clark et al. (2022), who used 25-km grid reference simulations in multiple climates to this end.

We show results from multiple 360-day-long coarse-grid prognostic runs using four randomly seeded pairs of NN models for temperature and humidity tendencies and radiative surface fluxes as well as an ensemble of the four pairs of NNs. Bias and RMSE are reported with respect to the fine-grid X-SHiELD reference model. We observe robust improvements across all NNs tested in time-mean land surface precipitation (9–25% lower RMSE) and land surface temperature (6–25% lower RMSE) with respect to a non-ML-corrected baseline simulation. Seasonally averaged land surface precipitation RMSE is also robustly improved across all seasons for all ML-corrected runs. Seasonally averaged land surface temperature RMSE is consistently improved during boreal summer across ML models, and two of four also reduce the boreal winter cold biases at high northern latitudes.

All ML-corrected simulations ran to completion without any crashes or runaway drifts. We tested our best-performing model (seed 0) over a five year simulation. It maintains a stable climate throughout the run, with consistent seasonal cycles year over year that closely repeat its first-year behavior.

Like the baseline model, the ML-corrected prognostic simulations develop significant seasonal biases in precipitable water, precipitation and vertical motion. The ML models somewhat overcorrect many of the tropical circulation biases of the baseline model.

Tropical precipitation variability is significantly reduced in the ML-corrected runs. While the baseline run produces too much extreme precipitation, the ML-corrected runs produce too little. Both the baseline and ML-corrected runs have weaker MJO and Kelvin

wave propagation than the fine-grid reference, with the ML-corrected run having the weakest wave propagation of the three.

Our results are encouraging for the prospect that ML trained on GSRM simulations can improve coarser-grid climate models. To realize this prospect, future work is still needed to improve the climate drift, circulation biases, and precipitation variability of ML-corrected runs.

7 Open Research

The code and experiment configurations needed to reproduce this work are available at the Github repository <https://github.com/ai2cm/nudge-to-3km-PIRE-1yr-workflow> which is archived at Zenodo (<https://doi.org/10.5281/zenodo.7063087>). The coarsened fine-grid data used for initial conditions and in the nudged coarse-grid simulation is available upon request through a Google Cloud Storage ‘requester pays’ bucket. The ERA5 data was obtained at <https://doi.org/10.24381/cds.f17050d7>. The GPCP data was obtained at <https://psl.noaa.gov/data/gridded/data.gpcp.html>.

Acknowledgments

We thank the Allen Institute for Artificial Intelligence for supporting this work and NOAA-GFDL for running the 1-year X-SHIELD simulation on which our ML is trained using the Gaea computing system. We also acknowledge NOAA-GFDL, NOAA-EMC, and the UFS community for making code and software packages publicly available.

References

- Adler, R. F., Sapiiano, M. R. P., Huffman, G. J., Wang, J.-J., Gu, G., Bolvin, D., ... Shin, D.-B. (2018). The global precipitation climatology project (gpcp) monthly analysis (new version 2.3) and a review of 2017 global precipitation. *Atmosphere*, 9(4). Retrieved from <https://www.mdpi.com/2073-4433/9/4/138> doi: 10.3390/atmos9040138
- Biewald, L. (2020). *Experiment tracking with weights and biases*. Retrieved from <https://www.wandb.com/> (Software available from wandb.com)
- Brenowitz, N. D., & Bretherton, C. S. (2019). Spatially extended tests of a neural network parametrization trained by coarse-graining. *J. Adv. Model. Earth Syst.*,

- 11, 2728-2744. doi: 10.1029/2019MS001711
- Bretherton, C. S., Henn, B., Kwa, A., Brenowitz, N. D., Watt-Meyer, O., McGibbon, J., ... Harris, L. (2022). Correcting coarse-grid weather and climate models by machine learning from global storm-resolving simulations. *Journal of Advances in Modeling Earth Systems*, 14(2), e2021MS002794. Retrieved from <https://agupubs.onlinelibrary.wiley.com/doi/abs/10.1029/2021MS002794> doi: <https://doi.org/10.1029/2021MS002794>
- Chen, G., Held, I. M., & Robinson, W. A. (2007). Sensitivity of the latitude of the surface westerlies to surface friction. *Journal of the Atmospheric Sciences*, 64(8), 2899 - 2915. Retrieved from <https://journals.ametsoc.org/view/journals/atasc/64/8/jas3995.1.xml> doi: 10.1175/JAS3995.1
- Cheng, K.-Y., Harris, L., Bretherton, C., Merlis, T. M., Bolot, M., Zhou, L., ... Fueglistaler, S. (2022). Impact of warmer sea surface temperature on the global pattern of intense convection: Insights from a global storm resolving model. *Geophysical Research Letters*, 49(16), e2022GL099796. Retrieved from <https://agupubs.onlinelibrary.wiley.com/doi/abs/10.1029/2022GL099796> (e2022GL099796 2022GL099796) doi: <https://doi.org/10.1029/2022GL099796>
- Clark, S. K., Brenowitz, N. D., Henn, B., Kwa, A., McGibbon, J., Perkins, W. A., ... Harris, L. M. (2022). Correcting a 200-km resolution climate model in multiple climates by machine learning from 25-km resolution simulations. *Journal of Advances in Modeling Earth Systems*, e2022MS003219. Retrieved from <https://agupubs.onlinelibrary.wiley.com/doi/abs/10.1029/2022MS003219> (e2022MS003219 2022MS003219) doi: <https://doi.org/10.1029/2022MS003219>
- Hersbach, H., Bell, B., Berrisford, P., Biavati, G., Horányi, A., Muñoz-Sabater, A., ... Thépaut, J.-N. (2019). *ERA5 monthly averaged data on single levels from 1979 to present*. Copernicus Climate Change Service (C3S) Climate Data Store (CDS).
- Li, L., Jamieson, K., DeSalvo, G., Rostamizadeh, A., & Talwalkar, A. (2018). Hyperband: A novel bandit-based approach to hyperparameter optimization. *Journal of Machine Learning Research*, 18(185), 1–52. Retrieved from <http://jmlr.org/papers/v18/16-558.html>
- McGibbon, J., Brenowitz, N. D., Cheeseman, M., Clark, S. K., Dahm, J., Davis, E., ... Fuhrer, O. (2021). fv3gfs-wrapper: a python wrapper of the FV3GFS

- 541 atmospheric model. *Geosci. Model Dev. Disc.*, 14, 4401–4409. doi: 10.5194/
542 gmd-14-4401-2021
- 543 Stevens, B., Satoh, M., Auger, L., Biercamp, J., Bretherton, C. S., Chen, X., ...
544 Zhou, L. (2019). DYAMOND: the dynamics of the atmospheric general circula-
545 tion modeled on non-hydrostatic domains. *Prog. Earth Planet. Sci.*, 6, 61. doi:
546 10.1186/s40645-019-0304-z
- 547 Watt-Meyer, O., Brenowitz, N. D., Clark, S. K., Henn, B., Kwa, A., McGibbon, J.,
548 ... Bretherton, C. S. (2021). Correcting weather and climate models by machine
549 learning nudged historical simulations. *Geophys. Res. Lett.*, 48, e2021GL092555.
550 doi: 10.1029/2021GL092555
- 551 Woelfle, M. D., Yu, S., Bretherton, C. S., & Pritchard, M. S. (2018). Sensitivity
552 of coupled tropical pacific model biases to convective parameterization in cesm1.
553 *Journal of Advances in Modeling Earth Systems*, 10(1), 126-144. Retrieved from
554 <https://agupubs.onlinelibrary.wiley.com/doi/abs/10.1002/2017MS001176>
555 doi: <https://doi.org/10.1002/2017MS001176>
- 556 Yanai, M., Esbensen, S., & Chu, J.-H. (1973). Determination of bulk properties of
557 tropical cloud clusters from large-scale heat and moisture budgets. *J. Atmos. Sci.*,
558 30, 611 - 627. doi: 10.1175/1520-0469(1973)030<0611:DOBPOT>2.0.CO;2
- 559 Yuval, J., & O’Gorman, P. (2020). Stable machine-learning parameterization of sub-
560 grid processes for climate modeling at a range of resolutions. *Nat. Commun.*, 11,
561 3295. doi: 10.1038/s41467-020-17142-3
- 562 Zhao, M. (2014). An investigation of the connections among convection, clouds,
563 and climate sensitivity in a global climate model. *Journal of Climate*, 27(5), 1845
564 - 1862. Retrieved from [https://journals.ametsoc.org/view/journals/clim/](https://journals.ametsoc.org/view/journals/clim/27/5/jcli-d-13-00145.1.xml)
565 [27/5/jcli-d-13-00145.1.xml](https://journals.ametsoc.org/view/journals/clim/27/5/jcli-d-13-00145.1.xml) doi: 10.1175/JCLI-D-13-00145.1
- 566 Zhou, L., Harris, L., & Chen, J.-H. (2022). *The gfdl cloud microphysics parameteri-*
567 *zation* (Tech. Rep.). Geophysical Fluid Dynamics Laboratory (U.S.). doi: [https://](https://doi.org/10.25923/pz3c-8b96)
568 doi.org/10.25923/pz3c-8b96
- 569 Zhou, L., Lin, S.-J., Chen, J.-H., Harris, L. M., Chen, X., & Rees, S. L. (2019).
570 Toward convective-scale prediction within the next generation global prediction
571 system. *Bulletin of the American Meteorological Society*, 100(7), 1225 - 1243.
572 Retrieved from [https://journals.ametsoc.org/view/journals/bams/100/7/](https://journals.ametsoc.org/view/journals/bams/100/7/bams-d-17-0246.1.xml)
573 [bams-d-17-0246.1.xml](https://journals.ametsoc.org/view/journals/bams/100/7/bams-d-17-0246.1.xml) doi: 10.1175/BAMS-D-17-0246.1