

Postgraduate Certificate in AI in Weather Prediction

Deep Learning Techniques For Prediction

Deep learning has become a cornerstone of modern weather prediction, offering powerful tools to extract patterns from massive, heterogeneous datasets. In the context of a postgraduate certificate in AI for weather prediction, mastering the vocabulary associated with deep learning techniques is essential for both theoretical understanding and practical implementation. This document provides an extensive explanation of the most important terms, illustrated with examples that relate directly to atmospheric science, and discusses common challenges that practitioners encounter.

Neural network refers to a computational model inspired by the biological brain, composed of layers of interconnected neurons. Each neuron receives a set of inputs, applies a weighted sum, adds a bias term, and passes the result through an activation function. In weather forecasting, a simple feed-forward neural network might be used to predict daily maximum temperature from a vector of historical observations, while more sophisticated architectures handle spatial and temporal dimensions.

Deep learning distinguishes itself from shallow learning by employing many hidden layers, enabling the model to learn hierarchical representations. For instance, a deep convolutional network can first detect edges in satellite imagery, then combine edges into cloud patterns, and finally infer precipitation intensity. The depth of the network often correlates with its ability to capture complex atmospheric processes, but it also raises concerns about over-parameterization and training stability.

Activation function introduces non-linearity, allowing the network to model intricate relationships beyond simple linear regression. Common choices include the Rectified Linear Unit (ReLU), the sigmoid function, and the hyperbolic tangent (tanh). In precipitation nowcasting, ReLU is favored for its computational efficiency and reduced likelihood of vanishing gradients, whereas sigmoid may be employed in the final layer of a binary classification model that predicts the occurrence of thunderstorms.

Loss function quantifies the discrepancy between the model's predictions and the true observations. Selecting an appropriate loss is critical for weather applications. A mean squared error (MSE) loss is typical for regression tasks such as temperature forecasting, while a cross-entropy loss is used for classification problems like distinguishing convective from stratiform rain. For probabilistic forecasts, the continuous ranked probability score (CRPS) provides a proper scoring rule that evaluates the entire predictive distribution.

Backpropagation is the algorithmic backbone that computes the gradient of the loss with respect to each model parameter. By propagating errors backward through the network, backpropagation enables the adjustment of weights via an optimization routine. In practice, backpropagation is implemented automatically by deep learning libraries, but understanding its mechanics helps diagnose issues such as

exploding gradients that can destabilize training on long time-series data.

Gradient descent describes the iterative process of updating model parameters in the direction that most rapidly reduces the loss. The simplest form, known as batch gradient descent, computes gradients using the entire training dataset. However, weather datasets are often too large to fit in memory, prompting the use of stochastic or mini-batch variants.

Stochastic gradient descent (SGD) approximates the true gradient by sampling a subset of the data at each iteration. This introduces noise that can help escape shallow local minima, a useful property when training deep networks on highly non-convex loss surfaces. In operational forecasting pipelines, SGD enables rapid updates as new observations become available, supporting continual learning.

Batch size denotes the number of training samples processed before a parameter update. Smaller batches increase the stochasticity of the gradient estimate, which can improve generalization but also lead to noisy convergence. For high-resolution radar data, a batch size of 16 or 32 may balance GPU memory constraints with training stability.

Epoch represents one full pass through the entire training dataset. Monitoring loss across epochs helps detect over-fitting, where the model's performance improves on the training set but degrades on unseen validation data. In the context of seasonal forecasting, a typical training regime might involve 50–100 epochs, with early stopping criteria to prevent unnecessary computation.

Over-fitting occurs when a model captures noise or idiosyncrasies of the training data rather than the underlying atmospheric dynamics. Symptoms include a widening gap between training and validation loss. Techniques such as regularization, dropout, and data augmentation are employed to mitigate over-fitting, especially when the historical record is limited.

Under-fitting describes a model that is too simple to capture the complexity of weather phenomena, leading to high bias and poor performance on both training and validation sets. Increasing network depth, adding more features (e.g., Humidity, wind shear), or employing richer architectures can address under-fitting.

Regularization adds a penalty term to the loss function to discourage overly complex models. Common forms include L1 regularization, which promotes sparsity by penalizing the absolute value of weights, and L2 regularization, which penalizes the squared magnitude. In temperature forecasting, L2 regularization can prevent weight explosion while preserving smooth spatial patterns.

Dropout randomly disables a fraction of neurons during each training iteration, forcing the network to develop redundant representations. A dropout rate of 0.2–0.5 is typical for fully connected layers in weather models. At inference time, dropout is disabled, and the full network is used for deterministic predictions. However, by keeping dropout active during inference and sampling multiple forward passes, practitioners can approximate Bayesian uncertainty (Monte Carlo dropout).

Early stopping monitors validation loss and halts training when improvement stalls for a predefined number of epochs. This prevents unnecessary epochs that could lead to over-fitting and reduces computational cost. In real-time forecasting systems, early stopping can also limit latency, ensuring that model updates finish within operational time windows.

Convolutional neural network (CNN) is a specialized architecture designed to process grid-structured data such as images. Convolutional layers apply learnable kernels that slide across the input, detecting local patterns regardless of their absolute position. In weather prediction, CNNs are widely used for precipitation nowcasting from satellite infrared imagery, where the spatial correlation of cloud structures is crucial.

A typical CNN architecture for rainfall estimation might begin with a series of convolution-pooling blocks that reduce spatial resolution while increasing feature depth, followed by a deconvolution (or transposed convolution) stage that restores the original resolution for pixel-wise prediction. This encoder-decoder design mirrors the popular U-Net architecture, which excels at semantic segmentation tasks such as identifying convective cores in radar fields.

Recurrent neural network (RNN) processes sequential data by maintaining a hidden state that evolves over time. The hidden state enables the network to capture temporal dependencies, a valuable property for forecasting variables that exhibit autocorrelation, such as temperature or humidity. However, standard RNNs suffer from vanishing or exploding gradients when modeling long sequences, limiting their practical horizon.

Long short-term memory (LSTM) and Gated recurrent unit (GRU) are gated variants of RNNs that address gradient issues by controlling information flow with input, forget, and output gates (LSTM) or reset and update gates (GRU). LSTMs have been successfully applied to multi-step temperature forecasts, where the model ingests a sequence of past observations and produces a forecast horizon of several days. GRUs, being computationally lighter, are attractive for applications with limited GPU resources, such as on-board processing for unmanned aerial vehicles collecting atmospheric data.

Attention mechanism allows a model to weigh different parts of the input sequence when generating each output element. By learning attention scores, the network can focus on the most relevant time steps or spatial locations. In weather prediction, attention can be used to highlight the influence of recent convective events on future precipitation, improving interpretability and forecast skill.

Transformer architecture replaces recurrent connections with self-attention layers, enabling parallel processing of the entire sequence. Transformers have demonstrated state-of-the-art performance in language modeling and are increasingly adopted for spatiotemporal forecasting. A weather-specific transformer might ingest multi-modal inputs—satellite images, radar fields, and numerical model outputs—encode them with positional embeddings, and decode a sequence of future forecasts. Because transformers scale quadratically with sequence length, careful design (e.g., Using sparse attention) is required for high-resolution, long-range forecasting.

Encoder-decoder models consist of an encoder that compresses input information into a latent representation, and a decoder that expands this representation into the desired output. In rainfall nowcasting, the encoder may be a CNN that extracts cloud features from the latest radar frame, while the decoder predicts the evolution of reflectivity over the next hour. The encoder-decoder paradigm is also employed in sequence-to-sequence translation of weather codes, such as converting raw model output into human-readable warnings.

Time series forecasting involves predicting future values of a variable based on its past observations. Deep learning models for time series can be purely recurrent, purely convolutional (temporal convolutions), or hybrid (e.G., ConvLSTM). A ConvLSTM combines convolutional operations with LSTM gates, preserving spatial structure while modeling temporal dynamics. This hybrid is particularly effective for predicting the movement of mesoscale convective systems across a radar mosaic.

Spatiotemporal modeling captures both the spatial and temporal dimensions of atmospheric phenomena. Techniques such as 3-D convolutions, ConvLSTM, and graph neural networks (GNNs) are used to model the evolution of variables across latitude, longitude, altitude, and time. For example, a 3-D CNN might predict the vertical profile of temperature from a sequence of satellite sounder observations, while a GNN could represent weather stations as nodes linked by distance and wind flow, enabling the model to learn localized interactions.

Data assimilation is the process of integrating observations into a numerical weather prediction (NWP) model to produce a more accurate initial state. Deep learning can augment assimilation by learning observation operators that map raw sensor data to model variables, or by providing rapid surrogate models that approximate the expensive assimilation step. A neural network trained to correct biases in satellite radiance retrievals can improve the quality of the assimilated state, leading to better forecasts.

Satellite imagery provides frequent, global coverage of cloud systems, sea surface temperature, and land surface characteristics. Convolutional networks are the natural choice for extracting features from these images. For instance, a CNN trained on infrared channels can classify cloud types, while a multitask network can simultaneously predict cloud top temperature and precipitation probability. Transfer learning from pretrained image classification models (e.G., ResNet) accelerates training and improves performance when labeled data are scarce.

Radar reflectivity measures the returned power from precipitation particles and is a key variable for short-range forecasting. Deep learning models often ingest reflectivity fields as multi-channel images, where each channel corresponds to a different elevation angle. A U-Net can be employed to segment convective cores, and a separate regression head can estimate rain rate. By incorporating motion vectors derived from optical flow, the model can predict the trajectory of storms over the next few hours.

Numerical weather prediction (NWP) solves the governing equations of fluid dynamics on a discretized grid. While NWP provides a physically based forecast, it is computationally intensive and can suffer from

systematic errors. Deep learning can be used to post-process NWP output, correcting biases (bias correction), downscaling (super-resolution), or generating probabilistic ensembles. A hybrid system might feed NWP fields into a CNN that learns the mapping from coarse-resolution forecasts to high-resolution precipitation estimates.

Ensemble methods generate multiple forecasts to quantify uncertainty. In deep learning, ensembles can be created by training several models with different random initializations, by varying hyperparameters, or by using dropout at inference time. The spread of the ensemble provides a measure of forecast confidence. For example, an ensemble of ConvLSTM models can produce a probabilistic precipitation forecast, where each member predicts a possible rainfall distribution.

Model bias refers to systematic errors that cause forecasts to consistently over- or under-predict a variable. Bias may stem from training data imbalances, sensor errors, or model architecture limitations. Detecting bias involves statistical analysis of residuals across seasons, geographic regions, and meteorological regimes. Correcting bias can be achieved through post-processing techniques such as quantile mapping, or by incorporating bias-aware loss functions during training.

Variance captures the sensitivity of the model to fluctuations in the training data. High variance models are prone to over-fitting. Reducing variance can be accomplished by increasing training data, simplifying the network, or applying regularization methods like dropout. In weather prediction, where training data may be limited for rare extreme events, controlling variance is essential to avoid spurious alarms.

Hyperparameter tuning involves selecting values for parameters that govern model training but are not learned directly, such as learning rate, batch size, and number of layers. Systematic search strategies include grid search, random search, and Bayesian optimization. For weather applications, hyperparameter tuning must consider operational constraints: A model that achieves marginally better skill but requires twice the inference time may be impractical for real-time alerts.

Learning rate determines the step size taken during gradient descent. Too large a learning rate can cause divergence, while too small a rate leads to slow convergence. Adaptive optimizers such as Adam automatically adjust the learning rate for each parameter based on first- and second-order moments of the gradients, often delivering faster training for deep weather models. Nevertheless, a manual learning-rate schedule (e.g., Cosine annealing) can further improve performance.

Optimizer is the algorithm that updates model weights based on computed gradients. Aside from SGD and Adam, other optimizers include RMSprop, AdaGrad, and Nadam. Each optimizer has distinct characteristics: RMSprop adapts learning rates based on a moving average of squared gradients, making it suitable for non-stationary weather data where gradient magnitudes can vary dramatically.

Weight initialization sets the initial values of network parameters before training begins. Proper initialization helps avoid vanishing or exploding activations. Techniques such as Xavier (Glorot) initialization are designed

for layers with symmetric activation functions, while He initialization is geared toward ReLU-based networks. In practice, initializing a CNN for satellite cloud classification with He initialization often yields faster convergence.

Feature scaling normalizes input variables to a common range, improving numerical stability and training speed. Common approaches include min-max scaling, which maps values to $[0, 1]$, and standardization, which subtracts the mean and divides by the standard deviation. For multi-modal weather data, each modality may require a different scaling strategy—for example, temperature fields are standardized, while radar reflectivity may be log-transformed before scaling.

Data augmentation artificially expands the training set by applying transformations that preserve the underlying physical meaning. For image-based weather data, augmentations include rotations, flips, and random cropping. Temporal augmentations, such as time-shifting or adding synthetic noise, can increase robustness to sensor errors. Augmentation is especially valuable when training deep models on limited severe-storm datasets.

Transfer learning leverages knowledge acquired from a source task to accelerate learning on a target task. In weather prediction, a model pretrained on a large generic image dataset (e.g., ImageNet) can be fine-tuned on satellite cloud images, reducing the amount of labeled data required. Transfer learning also enables the reuse of climate-scale representations when moving from global to regional forecasting.

Pretraining involves training a model on a large, often unsupervised, dataset before fine-tuning on a specific downstream task. Self-supervised objectives such as contrastive learning or masked image modeling can be applied to vast archives of satellite imagery, allowing the network to learn atmospheric textures without explicit labels. Subsequent fine-tuning on a labeled precipitation dataset yields improved performance compared to training from scratch.

Loss landscape visualizes how the loss function varies with respect to model parameters. Understanding the geometry of the loss landscape helps explain why certain optimizers converge faster or become trapped in local minima. In deep weather models, flat minima are associated with better generalization, while sharp minima may indicate over-fitting to specific atmospheric regimes.

Gradient vanishing occurs when gradients become extremely small as they propagate backward through many layers, hindering effective learning. This problem is mitigated by using ReLU activations, appropriate weight initialization, and architectural features such as residual connections. Conversely, exploding gradients can be controlled by gradient clipping, which caps the norm of gradients to a predefined threshold.

Batch normalization normalizes layer inputs across a mini-batch, stabilizing training and allowing higher learning rates. It also acts as a regularizer. In weather models, batch normalization is often applied after convolutional layers processing radar fields, improving convergence when training on heterogeneous

datasets that span multiple seasons.

Layer normalization performs a similar normalization but across feature dimensions within each sample, making it more suitable for recurrent architectures where batch sizes may be small. Layer normalization is commonly used in transformer models for atmospheric time-series forecasting, ensuring consistent scaling of hidden states across time steps.

Residual connections (or skip connections) add the input of a layer to its output, facilitating gradient flow and enabling the training of very deep networks. The ResNet architecture introduced this concept and has been adapted for weather prediction tasks such as super-resolution of NWP fields. By allowing low-frequency information to bypass several layers, residual connections help preserve large-scale atmospheric features.

Skip connections in encoder-decoder networks transmit high-resolution features from the encoder directly to the decoder, improving detail recovery. The U-Net architecture, which combines skip connections with a symmetric encoder-decoder, is widely used for pixel-wise segmentation of radar reflectivity, enabling accurate delineation of convective cores while maintaining context.

ResNet (Residual Network) typically consists of stacked residual blocks, each containing two or three convolutional layers with batch normalization and ReLU activation. In climate downscaling, a ResNet can learn the mapping from coarse global climate model outputs to high-resolution regional temperature fields, preserving large-scale patterns while adding fine-grained variability.

U-Net is a specific encoder-decoder architecture designed for segmentation. Its contracting path captures context, while the expansive path enables precise localization. For rainfall estimation, a U-Net can predict a dense rainfall map from a stack of satellite channels, with the skip connections ensuring that small cloud features are not lost.

Segmentation partitions an input image into meaningful regions. In weather, segmentation might separate rain-producing clouds from non-rain clouds, or delineate hail cores within a radar volume. Segmentation models output a probability map for each class, allowing thresholding to generate binary masks for downstream decision-making.

Classification assigns a single label to an input, such as predicting whether a given radar image contains a tornado-producing supercell. Multi-label classification extends this to scenarios where multiple phenomena may coexist, for instance, identifying both hail and strong wind signatures in the same radar scan.

Regression predicts continuous values, such as the amount of accumulated precipitation. Deep regression models often use a linear activation in the final layer to produce unrestricted outputs, though domain-specific constraints (e.g., Non-negative rainfall) can be enforced by applying a ReLU or exponential activation.

Probabilistic forecasting provides a full predictive distribution rather than a single deterministic value. Techniques include quantile regression, where the model learns to predict specific quantiles (e.G., 10th, 50th, 90th percentiles) of the distribution, and Bayesian neural networks, which treat weights as probability distributions. Probabilistic forecasts enable users to assess risk and make informed decisions under uncertainty.

Quantile regression minimizes the quantile loss, allowing the model to learn asymmetric error penalties. For precipitation, predicting the 95th percentile can be valuable for extreme-event warnings, while the median forecast offers a central tendency estimate. By training a single network to output multiple quantiles, forecasters obtain a calibrated prediction interval.

Bayesian neural network (BNN) places a prior distribution over network weights and updates this distribution using observed data, yielding a posterior that captures epistemic uncertainty. Exact inference is intractable for deep networks, so approximations such as variational inference or Monte Carlo dropout are used. BNNs are attractive for forecasting rare severe weather, where uncertainty quantification is crucial.

Monte Carlo dropout approximates Bayesian inference by performing multiple stochastic forward passes with dropout active, and aggregating the results to estimate predictive mean and variance. This technique is computationally cheap and integrates seamlessly with existing architectures, making it a practical choice for operational precipitation nowcasting where rapid uncertainty estimates are required.

Uncertainty quantification encompasses methods for estimating both aleatoric (inherent) and epistemic (model) uncertainty. Aleatoric uncertainty can be modeled by predicting a variance term alongside the mean, while epistemic uncertainty is captured by ensembles or Bayesian approaches. In weather prediction, separating these uncertainties helps identify whether forecast errors stem from observational noise (e.G., Radar attenuation) or from model limitations.

Calibration evaluates whether predicted probabilities correspond to observed frequencies. A well-calibrated model for severe thunderstorm warnings would produce a 70% probability that, over many events, roughly 70% actually materialize. Calibration can be assessed using reliability diagrams and improved through techniques such as isotonic regression or temperature scaling.

Reliability diagram plots observed event frequencies against predicted probabilities, revealing systematic miscalibration. In practice, a reliability diagram for a deep learning hail-probability model may show over-confidence at high probability levels, prompting post-processing adjustments.

Skill scores quantify forecast performance relative to a reference, such as climatology or persistence. Common skill scores include the Brier score for probabilistic binary events, the continuous ranked probability score (CRPS) for continuous variables, and the equitable threat score (ETS) for categorical events. Deep learning models are typically evaluated against these benchmarks to demonstrate added value.

Brier score measures the mean squared difference between predicted probabilities and binary outcomes.

Lower Brier scores indicate better calibrated and more accurate predictions. For a deep model that predicts the probability of flash flood occurrence, the Brier score provides a single metric that combines calibration and discrimination.

CRPS generalizes the Brier score to continuous variables by integrating the squared difference between the cumulative forecast distribution and the observed value. CRPS is widely used to evaluate probabilistic temperature or precipitation forecasts generated by deep ensembles or quantile regression networks.

Operational forecasting refers to the deployment of models in real-time environments where timeliness, reliability, and interpretability are paramount. Deep learning models must be integrated into existing pipelines, often requiring conversion to optimized formats (e.G., ONNX), and must meet strict latency constraints to support early warning issuance.

Real-time inference is the process of generating predictions as new observations become available. For severe weather, inference must be completed within minutes to be actionable. Techniques such as model quantization, pruning, and GPU acceleration are employed to reduce latency while preserving forecast skill.

GPU acceleration leverages graphics processing units to parallelize the large matrix operations inherent in deep learning. Modern weather models often run on clusters equipped with NVIDIA or AMD GPUs, enabling the training of multi-gigabyte networks on billions of observations. For inference, a single GPU can process thousands of radar frames per second, supporting high-frequency nowcasting.

Parallel computing distributes computation across multiple processors or nodes. Data parallelism replicates the model on each GPU and synchronizes gradients, while model parallelism splits the network layers across devices. In large-scale climate downscaling, data parallelism is commonly used to accelerate training on distributed clusters.

Distributed training extends parallelism across a network of machines, employing frameworks such as Horovod or PyTorch Distributed. Proper scaling requires careful handling of communication overhead, gradient aggregation, and checkpointing. For weather prediction, distributed training enables the use of massive datasets that span decades of reanalysis and satellite archives.

Model interpretability addresses the need to understand how a deep learning model arrives at its predictions. Techniques such as SHAP values, saliency maps, and class activation mapping reveal which input features or regions most influence the output. In the context of tornado prediction, interpretability can highlight the atmospheric layers that the model deems most critical, fostering trust among meteorologists.

SHAP values (SHapley Additive exPlanations) assign an importance score to each input feature based on cooperative game theory. By computing SHAP values for a temperature forecast model, analysts can quantify the contribution of humidity, wind shear, and sea-surface temperature to each prediction, enabling transparent decision support.

Saliency maps visualize the gradient of the output with respect to the input, indicating which pixels in a satellite image most affect the forecast. A saliency map for a convective initiation model may highlight bright infrared regions where the network focuses its attention, providing insight into the physical basis of the prediction.

Explainable AI (XAI) encompasses a suite of methods and principles aimed at making AI systems more transparent, accountable, and trustworthy. In weather forecasting, XAI helps bridge the gap between data-driven models and domain expertise, ensuring that forecasts align with physical understanding and operational requirements.

Domain adaptation mitigates performance degradation when a model trained on one data distribution is applied to another. For example, a model trained on European radar data may need adaptation to work on North American networks with different beam configurations. Techniques such as adversarial training or feature alignment adjust the model to new domains without extensive retraining.

Climatology provides long-term statistical averages that serve as baselines for forecast verification. Deep learning models can incorporate climatological priors, for instance by bias-correcting NWP outputs toward historical means, thereby improving skill in regions with sparse observational coverage.

Climate change introduces non-stationarity into atmospheric data, challenging the assumption that past patterns will persist. Model retraining, continual learning, and incorporation of climate-scenario inputs are strategies to maintain forecast relevance as the climate evolves.

Extreme events such as hurricanes, tornadoes, and flash floods are rare but high-impact. Deep learning models often struggle with limited training examples, leading to poor generalization. Synthetic data generation, oversampling, and specialized loss functions (e.g., Focal loss) can improve detection of these low-frequency phenomena.

Severe weather encompasses phenomena that pose immediate threats to life and property. Real-time deep learning pipelines for severe weather must balance accuracy with speed, ensuring that false alarms are minimized while detection rates remain high. Operational metrics such as probability of detection (POD) and false alarm ratio (FAR) are used to evaluate system performance.

Tornado prediction is a challenging task due to the small spatial scale and short lead time of tornadoes. Convective-allowing models (CAMs) provide high-resolution inputs, and deep networks can learn to identify tornadic signatures from radar velocity fields. A common approach is to feed a sequence of velocity and reflectivity images into a ConvLSTM, outputting a tornado probability map.

Hail forecasting benefits from dual-polarization radar measurements that capture particle shape and size. Deep learning models can ingest the differential reflectivity (ZDR) and specific differential phase (KDP) channels to estimate hail size distribution. Quantile regression can then predict the 90th percentile hail size, informing aviation and public safety decisions.

Flood forecasting requires integrating precipitation forecasts with hydrological models. Deep learning can serve as a surrogate for the hydrological routing component, accelerating forecasts while preserving accuracy. A recurrent network trained on past rainfall-runoff pairs can predict river discharge in near-real time, enabling early flood warnings.

Precipitation nowcasting aims to predict rainfall over the next hour at high spatial resolution. State-of-the-art nowcasting systems combine optical flow methods with deep learning. A hybrid model may use a CNN to extract features from the current radar image, feed them into a ConvLSTM to model motion, and output a sequence of future reflectivity fields.

Data pipelines orchestrate the flow of raw observations through preprocessing, model inference, and post-processing stages. Robust pipelines handle heterogeneous data formats (e.G., NetCDF, GRIB), perform quality control, and ensure consistent timestamps. Automation tools such as Apache Airflow or Prefect are often employed to schedule and monitor pipeline execution.

Data preprocessing includes steps such as regridding, masking, and handling missing values. Regridding aligns data from different sources onto a common spatial grid, while masking removes regions with unreliable observations (e.G., Satellite view-angle extremes). Missing data can be imputed using temporal interpolation or learned generative models.

Missing data handling is critical because many atmospheric sensors produce gaps due to instrument outages or transmission errors. Simple approaches include forward-fill or linear interpolation, but advanced techniques use autoencoders to reconstruct missing fields based on spatial context, preserving physical consistency.

Outlier detection identifies anomalous observations that may corrupt training. Statistical methods (e.G., Z-score thresholds) or unsupervised deep models (e.G., Isolation forest) can flag extreme values. In radar, outliers often arise from ground clutter, which must be removed before model ingestion.

Time lag denotes the delay between the occurrence of a physical process and its observation. For satellite infrared channels, the time lag may be a few minutes, whereas ground-based observations can have larger delays. Accounting for time lag in model inputs improves alignment between cause and effect, enhancing forecast skill.

Autocorrelation measures the similarity of a time series with its lagged versions. Strong autocorrelation in temperature series justifies the use of lagged inputs in regression models. However, excessive reliance on autocorrelation can cause models to ignore exogenous drivers, reducing adaptability to sudden regime shifts.

Cross-validation assesses model generalization by partitioning the data into training and validation subsets. In time-series contexts, standard random splits can violate temporal dependencies, leading to optimistic estimates. Instead, a rolling-origin or forward-chaining cross-validation scheme respects chronological

order.

k-fold cross-validation divides the dataset into k equal parts, training on k – 1 folds and validating on the remaining fold. For weather data, a spatial k-fold approach may be used, where each fold corresponds to a geographic region, testing the model's ability to transfer across climate zones.

Rolling window validation trains the model on a sliding time window and validates on the subsequent period, mimicking operational forecasting. This method provides insight into how model performance evolves over time and can reveal degradation due to climate trends or sensor changes.

Hyperparameter search automates the exploration of hyperparameter space. Grid search exhaustively evaluates combinations, while random search samples randomly, often achieving comparable results with fewer evaluations. Bayesian optimization builds a probabilistic model of performance and selects promising hyperparameters, reducing the number of required experiments.

Grid search is straightforward but computationally expensive, especially when many hyperparameters (learning rate, batch size, dropout rate, number of layers) are considered. For deep weather models, a coarse grid followed by a refined random search is a practical compromise.

Random search samples hyperparameter configurations uniformly at random. Empirical studies have shown that random search can be more efficient than grid search because many hyperparameters have diminishing returns beyond certain thresholds. In practice, random search is often combined with early stopping to discard poorly performing trials early.

Bayesian optimization constructs a surrogate model (e.G., Gaussian process) of the objective function (validation loss) and selects hyperparameters that maximize expected improvement. This approach is well-suited for expensive training runs typical of large CNNs for satellite cloud classification.

AutoML (Automated Machine Learning) frameworks automate the entire pipeline from data preprocessing to model selection and hyperparameter tuning. In weather forecasting, AutoML can quickly generate baseline models, which experts then refine using domain knowledge.

Model deployment transitions a trained model into a production environment. Deployment considerations include model serialization (e.G., TorchScript, SavedModel), API design, latency budgeting, and resource allocation. A common pattern is to expose the model via a RESTful endpoint that receives input data (e.G., Radar fields) and returns forecast probabilities.

API integration connects the deep learning service with existing forecasting systems, such as the National Weather Service's AWIPS platform. Standard data exchange formats (e.G., JSON, Protocol Buffers) and authentication mechanisms ensure secure and reliable communication.

Containerization packages the model, its dependencies, and runtime environment into a portable unit.

Docker is the de-facto standard, allowing reproducible deployment across diverse hardware. By encapsulating the deep learning inference code in a container, operational teams can manage versioning and scaling more effectively.

Kubernetes orchestrates containers across a cluster, providing automated scaling, load balancing, and fault tolerance. In a large-scale nowcasting service, Kubernetes can spin up additional inference pods during peak demand (e.g., During a severe storm outbreak) and scale down when activity subsides.

Monitoring tracks model performance, system health, and data drift in real time. Metrics such as inference latency, GPU utilization, and forecast error statistics are logged and visualized. Alerts can be configured to trigger retraining when performance degrades beyond a threshold.

Drift detection identifies changes in the statistical properties of input data that may indicate sensor upgrades, new satellite platforms, or evolving climate patterns. Techniques include comparing feature distributions over sliding windows or using a dedicated drift detection model. Early detection enables timely model updates.

Model updating can be performed offline (periodic retraining) or online (continual learning). Offline updates are simpler to manage but may introduce latency in responding to new patterns. Online learning methods, such as incremental gradient updates, allow the model to adapt continuously while preserving previously learned knowledge.

Continuous learning addresses the catastrophic forgetting problem, where a model loses performance on earlier data when trained on new data.