

Supporting information for "SYNOPTICBENCH: Evaluating Vision-Language Models on Generating Weather Forecast Discussions of the Future"

June 2026

Contents of this File

- Implementation Details
- Prompt Details
- Additional Preprocessing Details
- Parameter Settings
- Data Quality Control
- Extraction Details
- LLM as a Judge
- Location Stop Nodes
- SPACE Temperature Scores
- Additional Example Cases
- Training Set Size

Implementation Details. We fine-tune each base model by training for roughly 1-3 epochs on one NVIDIA H200 GPU before running it on the test set. We use 2016-2022 for the training set, 2023 as the validation set, and 2024-2025 as the test set. The projector weights were fully fine-tuned, while Low-Rank Adaptation (LoRA) was used to train a small portion of the attention mechanism in the language model (<2%). We apply weight decay and dropout to the LoRA adapter branches to prevent overfitting. In addition, early stopping is used by computing SPACE-local scores for the validation set every 5000 steps.

Prompt Details. Base model and finetuned model versions for LLaVA-v1.5-7B, LLaVA-v1.5-13B, Qwen2-VL-7B, and LLaMA-3.2-11B-Vision were run on the test set. Each model uses the following prompt to generate discussions (see Fig. 1): "Analyze these weather charts showing maximum precipitation over the 3-day period (shaded contours), 500 mb geopotential height (unshaded contours) and 850 mb wind velocity (wind barbs) for the <[City, State]> forecast region and generate a detailed forecast discussion of the large-scale features relevant to the area. The forecast region is shown in the yellow box".

Additional Preprocessing Details. The images used in this study use the temporal mean of lead times ranging from 3 hours to 48 hours in 3 hour increments. The anomalies are calculated from monthly means ranging across the entire dataset. The month chose for climatology is always the same month that the AFD is issued. To create forecast images, 2-m temperature anomaly is capped between -5° and 5°C. 500 mb geopotential height contours are drawn at intervals of 60m and 850mb wind velocity barbs are plotted every 20th grid point to avoid visual clutter.

To filter AFDs, we only extract sentences containing the following keywords: *trough, trof, the low, this low, upper level low, low pressure, low-pressure, upper low, cyclone, closed low, cut-off low, troughing, ridge, the high, this high, upper level high, high pressure, high-pressure, upper high, anticyclone, blocking, ridging, cold front, warm front, warmer, cooler, and freezing*. We also exclude sentences with the following model keywords: *ECMWF, EURO, HRRR, ECCC, CMC, GEM, NAM, UK-MET, ICON, RAP, SREF, HREF*, and synoptic pressure system keywords: *shortwave, short-wave, sfc trough, surface trough, sfc ridge, surface ridge, surface low, sfc low, surface high, sfc high* to increase the relevancy of NWS text in the use case presented in this study.

Parameter Settings. The same settings for the finetuning parameters were used for LLaVA-v1.5-7B, LLaVA-v1.5-13B, Qwen2-VL-7B, and LLaMA-3.2-11B-Vision. The learning rate was set at $2e-5$, the batch size was set to 16, the LoRA rank was set to 8, and the LoRA alpha was set to 16. LoRA was applied to all linear modules within the attention and multi-layer perceptron blocks. AdamW was used as an optimizer. Image preprocessing was handled natively by each model’s respective Hugging Face AutoProcessor. The images were converted to RGB and dynamically resized, padded, and normalized according to the base model’s default vision encoder requirements (e.g., SigLIP for LLaVA/Llama and dynamic resolution for Qwen2.5-VL) before being tokenized.

Data Quality Control. Each text sample is matched to the closest GFS run. The initial raw dataset contained 1,367,041 samples. Text samples that are not within 12 hours of any GFS run are discarded. Pre-filtered samples in the raw dataset containing less than 30 words or greater than 200 words were also discarded. In addition, we removed samples containing text that suggested longer lead times beyond the 48-hour lead time post filtering (e.g. “Day 3”, “extended period”, “long term”). Quality control tests revealed zero samples containing text duplicates, unusual characters, absent image tokens, missing image files, or corrupted image files. Once all of the filtering specific to the example use case of this experiment was completed, the training set was reduced to 457,230 high-quality samples.

Extraction Details. We developed an automated Python pipeline to interface with

the IEM JSON API, systematically querying and downloading historical AFDs from 117 NWS Weather Forecast Offices. The raw text files were extracted for the period spanning January 2019 to November 2025, yielding a comprehensive collection of chronologically ordered NWS discussions. To ensure strict temporal accuracy for multimodal pairing, we implemented a custom regular expression parser to extract the exact issuance timestamp directly from the raw text body of each product, converting all times to Coordinated Universal Time (UTC).

LLM as a Judge. Gemini-2.5-Flash was used as a judge for evaluation. The following text was used as the prompt for the judge: "Compare the generated forecast discussion to the ground truth forecast discussion. Evaluate how accurately the generated discussion captured the synoptic-scale features. Ignore minor differences in phrasing and focus on meteorological accuracy and completeness. Score the prediction from 0 to 1:

- 0: Completely incorrect or severely hallucinated.
- .25: Major meteorological errors or missing critical synoptic features.
- .5: Partially correct, but missing some context or containing minor errors.
- .75: Mostly accurate and aligns well with the ground truth, minor omissions.
- 1: Excellent, highly accurate, and meteorologically sound compared to the ground truth.

Location Stop Nodes. The following locations are excluded from functioning as common relatives or matching to other locations via a common relative: "Canada", "CONUS", "Eastern Canada", "Central Canada", "Western Canada", "Eastern CONUS", "Western CONUS", "Central CONUS", "Central Plains", "Ohio Valley", "Great Lakes", "Central U.S.", "Eastern U.S.", "Western U.S.", "Central United States", "Eastern United States", "Western United States", "Central U", "Eastern U", "Western U", "Eastern US", "Western US", "Central US", "Midwest", "The Plains".

SPACE Temperature Scores. SPACE scores for temperature are shown in Table 4. The following key terms for temperature polarity were used: "warm front", "warmer than average", "warmer than normal", "above average temperature", "cold front", "cooler than average", "cooler than normal", "colder than normal", and "below average temperature". The number of temperature terms found in both AFDs and generated text was dominated by "cold front" and "warm front" (>80%). Fine-tuned Qwen2.5-VL-7B-LoRA had the highest SPACE-local scores while fine-tuned LLaVA-v1.5-7B had the highest SPACE-aggregate scores. The 48-hour temporal average that was used to create the images could lead to too much smoothing for the temporal scale at which fronts often take effect. Although SPACE scores could be effective for synoptic temperature systems, altering training images and experimental setup to target temperature-related phenomena would improve the reliability of evaluation in future work.

Additional Example Cases. Three additional example cases are shown in Figures A1-A3. In Case 3, there are two matching instances of pressure systems and two unmatching instances of pressure systems, leading to a coverage ratio of 0.5. Of the matching instances, both indicated high pressure, resulting in a match score of 1 and a SPACE-local score of 0.5. In Case 4, the prediction described synoptic pressure systems, while the reference text did not (Short wave troughs are not considered synoptic pressure systems). This results in a SPACE-local score that defaults to 0. Although there was some accurate information in the generated text, it was not relevant enough to the area to be included in the AFD. The pressure systems described in the generated

Model	s_s^{loc}	s_m^{loc}	r_c^{loc}	s_s^{agg}	s_m^{agg}	r_c^{agg}
LLaVA-v1.5-7B-base	.0173 $\pm$.0028	.8278 $\pm$.0526	.0207 $\pm$.0031	.3447 $\pm$.0028	.5602 $\pm$.0526	.6017 $\pm$.0031
LLaVA-v1.5-13B-base	.0793 $\pm$.0048	.9068 $\pm$.0125	.0862 $\pm$.0051	.4229 $\pm$.0008	.6053 $\pm$.0056	.6843 $\pm$.0013
LLaMA-3.2-11B-Vision-base	.0671 $\pm$.0025	.7294 $\pm$.0075	.0886 $\pm$.0030	.3300 $\pm$.0025	.5683 $\pm$.0075	.5709 $\pm$.0030
Qwen2.5-VL-7B-base	.0491 $\pm$.0032	.8330 $\pm$.0155	.0577 $\pm$.0035	.3302 $\pm$.0032	.5478 $\pm$.0155	.5878 $\pm$.0035
LLaVA-v1.5-7B-LoRA	.1863 $\pm$.0056	.8839 $\pm$.0075	.2077 $\pm$.0060	.4679 $\pm$.0056	.5965 $\pm$.0075	.7718 $\pm$.0060
LLaVA-v1.5-13B-LoRA	.1862 $\pm$.0063	.8697 $\pm$.0085	.2098 $\pm$.0067	.4209 $\pm$.0063	.5870 $\pm$.0085	.7043 $\pm$.0067
LLaMA-3.2-11B-Vision-LoRA	.1726 $\pm$.0055	.8992 $\pm$.0073	.1892 $\pm$.0057	.4597 $\pm$.0055	.6124 $\pm$.0073	.7395 $\pm$.0057
Qwen2.5-VL-7B-LoRA	.1886 $\pm$.0074	.8660 $\pm$.0058	.2128 $\pm$.0062	.3992 $\pm$.0074	.5978 $\pm$.0058	.6529 $\pm$.0062
Gemini-3.1-Pro	.0397 $\pm$.0050	.7725 $\pm$.0340	.0408 $\pm$.0057	.3963 $\pm$.0050	.5135 $\pm$.0340	.7732 $\pm$.0057
Nearest Neighbor	.0585 $\pm$.0036	.7532 $\pm$.0166	.0753 $\pm$.0041	.3216 $\pm$.0036	.5607 $\pm$.0166	.5614 $\pm$.0041
Climatology	.0850 $\pm$.0042	.7728 $\pm$.0148	.1073 $\pm$.0048	.4042 $\pm$.0042	.5795 $\pm$.0148	.6894 $\pm$.0048
Llama-3.2 (blind)	.0912 $\pm$.0048	.8664 $\pm$.0141	.1043 $\pm$.0051	.4211 $\pm$.0048	.5885 $\pm$.0141	.7050 $\pm$.0051

Table 1: SPACE was used for synoptic temperature system evaluation for base models, finetuned models, and baselines. We calculate the mean and standard error of the SPACE-local and SPACE-aggregate scores (s_s), match scores (s_m), and coverage ratios (r_c) within the test set. The range of possible scores for all metrics is 0 to 1, with 1 being a perfect score

text could still impact SPACE-aggregate scores. In Case 5, all pressure systems were described in the same general area, leading to a coverage ratio of 1. The match score was imperfect due to a ridge that would impact the region soon after the trough.

Training Set Size. The impact of the size of the training set on the SPACE scores is shown in Figure A4. The majority of the improvement in SPACE scores occurs within the first half epoch of training.

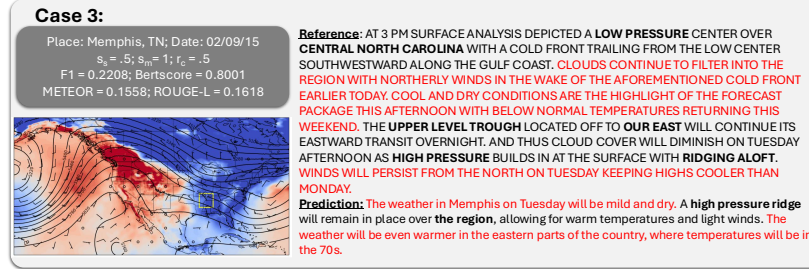


Figure A1: Example Case 3

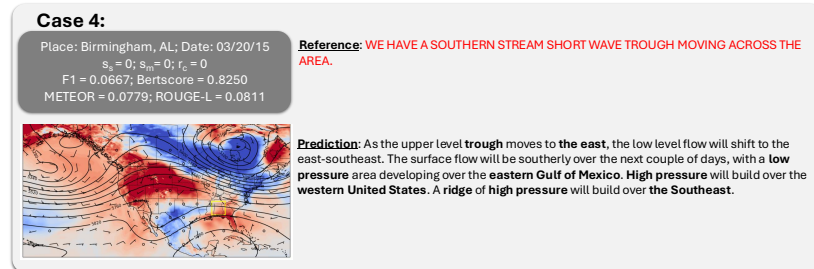


Figure A2: Example Case 4

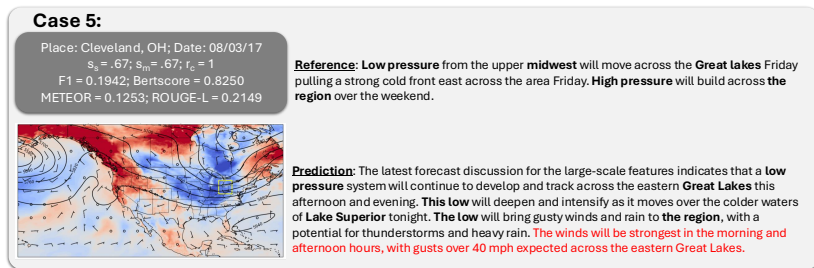


Figure A3: Example Case 5

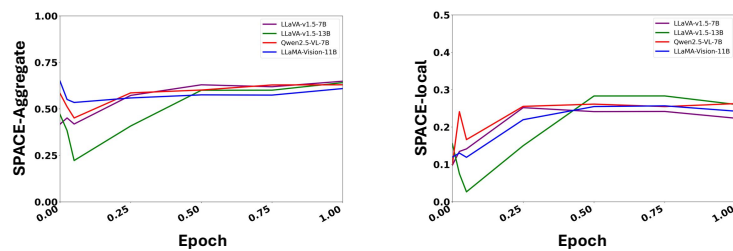


Figure A4: Training Set Size