

General statement

We would like to thank the editor for coordinating the review of our work and the peer-reviewers for their valuable comments on our study. In the following, we have addressed the referees' comments and presented our revisions in the manuscript. For clarity, our responses are highlighted in blue.

Referee comment #1

The manuscript describes and evaluates an ensemble nowcasting system, which has been constructed with use of a denoising diffusion probabilistic model (DDPM). Though diffusion models (DM) are recently intensively studied in meteorology and weather forecasting, their ability to simulate the wind uncertainty has got less attention. Thus, the presented results can be very informative for those who would like to build a complex probabilistic nowcasting system based on a diffusion model. The authors succeeded to demonstrate with both statistical evaluation and case studies that such a model could be used for nowcasting of wind speed. The DDPM itself and methods used for its evaluation are generally well described and justified, with some exceptions (which will be commented below). It must be also appreciated that the performance of several noise schedules (Linear, Cosine, Sigmoid) was thoroughly tested and analyzed. Still, there are several aspects where the model could be improved, above all concerning extreme events – as it is admitted in the conclusion of the manuscript. The authors claim several times that one of the biggest advantages of the DM-s against other AI/ML techniques is the maintenance of physical consistency. But the evaluation is done only with respect to their own analyses and reference forecasts, one cannot find comparisons of the DDPM outputs with performance of AI models based on different principles. Thus, these statements would need some clarification. Also, the presented case studies could be analyzed and discussed with more details.

The manuscript can be considered as an interesting pilot study, introducing the application of DDPM in wind nowcasting. Despite the above-mentioned weaknesses and some deficiencies in the presentation of the schemes and results, the manuscript exhibits an overall good quality and after corrections and clarifications, it could be appropriate for publication.

Reply:

We thank the reviewer for the positive and constructive assessment of our work. We have carefully addressed all the comments and revise the manuscript accordingly. In particular, we have removed the unsupported claims about physical consistency, added more details to the case studies, and clarified the limitations of the current study. The revisions are traceable in the revised manuscript.

Specific comments:

#Comments 1

Between the Lines 80-85: "The learning objective of DDPM is the error of SIVA nowcast, defined as the difference between forecast and the corresponding analysis field."

Reviewer (R): Usually, objective analysis also contains errors, above-all in grid-points laying far from observations, in complex terrain, etc. How it is in the case of SIVA nowcast, how much could the error of the analysis influence the DDPM training and the magnitude of forecast errors?

Can you estimate somehow the accuracy of analyses used (e.g. have you done cross-validation)?
Can you comment on this?

Reply 1:

Thanks for the comments. We agree that the analysis contains errors. To quantify its accuracy, we performed a cross-validation experiment using the same station data that go into the analysis as described in section 2. The stations were split into a training set (90% stations) used to produce the analysis, and a test set (10% stations) as out-of-samples. The cross-validation was performed separately for the training, validation, and test periods, with consistent results across all three. Table 1 shows the results for the test period (the period used to evaluate our DDPM results).

Table 1. Cross-validation scores of the SIVA analysis

Station set	train	test
bias (m s^{-1})	0.149	0.133
mae (m s^{-1})	0.351	0.725
rmse (m s^{-1})	0.454	0.943
Avg. wind (m s^{-1})	1.73	1.74

The results show that the analysis has a mean RMSE of about 0.94 m s^{-1} on out-of-sample stations. The spatial maps (Fig. 1) show that errors are a bit larger in the two mountainous areas ($30\sim31.5^\circ\text{N}$, $116\sim117^\circ\text{E}$ and $30\sim30.9^\circ\text{N}$, $117\sim120^\circ\text{E}$) in the southwest ($400\text{--}1200 \text{ m}$ elevation). Published results from the SIVA system (Zhu et al., 2025) and its older version INCA (Haiden et al., 2011) also support that the analysis is a reasonable approximation of the true state. Given this moderate error level, we treat the analysis as a practical approximation, acknowledging that it may introduce some uncertainty into our results.

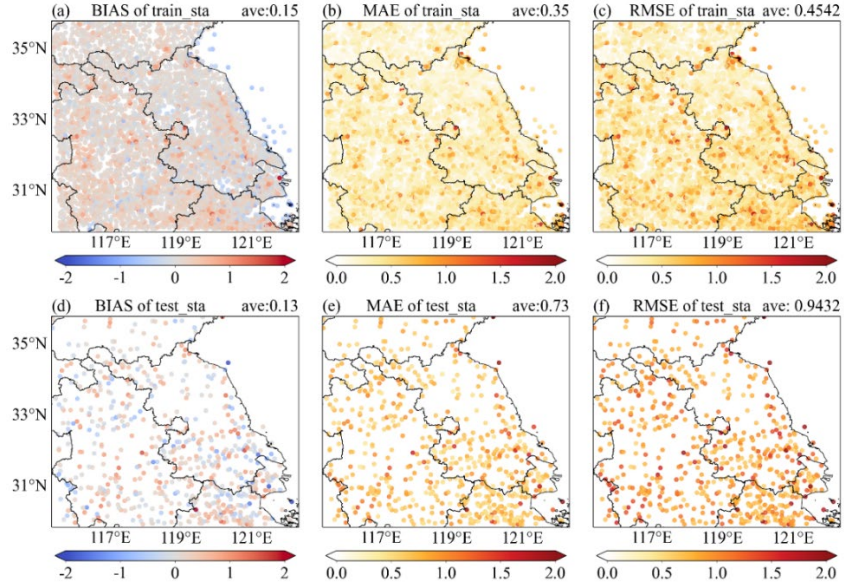


Fig.1. Spatial distribution of bias, MAE, and RMSE for the training stations (a–c) and test stations (d–f) in the test period. (a, d) Bias; (b, e) MAE; (c, f) RMSE.

Regarding the impact on DDPM training and the magnitude of forecast errors: we agree that errors in the analysis could affect both. The cross-validation results show that the analysis has relatively small errors (RMSE $\sim 0.94 \text{ m s}^{-1}$ on out-of-sample stations, and 0.45 m s^{-1} on training stations) with a mean wind speed of $\sim 1.74 \text{ m s}^{-1}$. Given this overall error level, we believe it is unlikely to overturn our main conclusions. Based on this moderate error level, we treat the

analysis as a gridded approximation of the true state.

We have not deeply investigated how the analysis errors propagate into the DDPM training or the final ensemble nowcasts. However, we fully agree with the comment that if one aims to represent forecast uncertainty against real observations, the influence of analysis errors should be considered. This is a limitation of the current study, and we have noted it as an important direction for future work, where independent station observations will be used to validate the results and to better understand the role of analysis errors.

We also added a short note in the Conclusions section i.e.: “In addition, the model is trained and evaluated using the analysis field as a reference, which itself contains errors. The impact of these errors on the training and verification has not been quantified, which constitutes a limitation. The absolute magnitude of the reported forecast errors may be affected, although the main conclusions, such as the relative performance of different noise schedules, are based on comparisons using the same reference. The ensemble nowcasts have also not been validated against independent station observations. Future work should address this limitation by using independent station observations to validate the results and to better understand the role of analysis errors.”, and the revision is in Line 436–441.

References:

Zhu, Y. W., Atencia, A., Dabernig, M., and Wang, Y.: Quantifying the analysis uncertainty for nowcasting application, *Geosci. Model Dev.*, 18, 1545–1559, 2025.

Haiden, T., Kann, A., Wittmann, C., Pistotnik, G., Bica, B., and Gruber, C.: The Integrated Nowcasting through Comprehensive Analysis (INCA) System and Its Validation over the Eastern Alpine Region, *Weather Forecast.*, 26, 166–183, 2011.

#Comments 2

85: The Fig. S1 in the supplementary materials shows the domain only very schematically. One has no information about the orography of the area (e.g. are there also high mountains?), which can be important from the wind nowcast point of view. Also, it is rather unusual that the domain is presented in the supplementary material and it is not one of the figures of the manuscript.

Reply 2:

Thanks for the suggestion. We agree that presenting the domain map only in the supplementary material is unusual and that orographic information is crucial for wind nowcasting. Accordingly, we have moved this figure to the main text as the new Figure 1 and added a subplot for topographic information. The revised figure and its caption are traceable in Line 95–96.

#Comments 3

85-90: „The data was split into nonoverlapping parts for training (1 October 2021–30 April 2023), validation (1 May 2023–31 May 2023), and testing (remainder).“

R: Later in the conclusion of the manuscript (lines 350-355) you mention that DDPM forecasts are still too smooth to reproduce some extreme events. Is it perhaps related to too short (~1,5 year) training period? Was there any particular reason for choosing the period of this length?

Reply 3:

Thanks for the comment. The training period (1 October 2021 – 30 April 2023) was determined

by the dataset that was accessible to us under the terms of our data agreement; a longer continuous period was not available at the time of this study. We agree that a longer training period would likely improve the representation of extreme events and reduce the excessive smoothness noted in the conclusion.

We therefore revised the conclusion to include the short training period as a contributing factor i.e.: “This limitation may stem from the use of a basic diffusion model architecture as well as the relatively short training period (~1.5 years), which likely does not provide sufficient samples of extreme wind events”. The revision is in Line 432–433.

#Comments 4

100-105: In the Section 3.1., the meaning of variables „q“ and „p“, which are products of Equations (1) and (2) is not explicitly mentioned, though one could guess that these are probability distributions. Also „I“ is not explained (Identity matrix?). I would recommend to specify it, for bigger clarity.

Reply 4:

Thanks for the suggestion. We have added explicit definitions of q , p , and I in Section 3.1 i.e.: “The $q(x_t|x_0)$ denotes the conditional probability distribution of x_t after adding t steps of Gaussian noise $\mathcal{N}(0,I)$ to x_0 , where I is the identity matrix” and “Here, $p_\theta(x_{t-1}|x_t)$ represents the reverse conditional probability distribution of recovering x_{t-1} from x_t , parameterized by a neural network θ .” The revisions are in Line 106–110.

Specifically, we stated that $q(x_t|x_0)$ is the conditional probability distribution of x_t after adding t steps of Gaussian noise to x_0 , with I as the identity matrix; and that $p_\theta(x_{t-1}|x_t)$ is the reverse conditional probability distribution for recovering x_{t-1} from x_t , parameterized by a neural network.

#Comments 5

Figure1: The figure might help the reader to understand the training process. But I do not find it to be self-explanatory enough and it takes some time to interpret it. Hence I would propose to improve this Figure and to relate it better with the corresponding text in 3.2. and equations in 3.1. For example, one should easily find, where is the start of the training (this is the creation of „Errors“ in the upper left corner of the Figure). You could highlight this step somehow (e.g. with a letter or number) and refer to it also in the text of 3.2. Similarly also some other parts of the DDPM framework, the cycle of the sampling process, etc. as one might understand in detail, what is depicted on the Figure. There is also a 3D graph (maybe a distribution function, product of the Denoiser?) on the middle, right hand side of the Figure, which has no name and explanation. Please, amend that as well.

Reply 5:

Thanks for this detailed suggestion. We agree that the current figure is not self-explanatory enough. We redesigned the figure as follows:

- Add numbered steps (e.g., ①, ②, ...) to clearly indicate the start of training (the creation of “Errors” in the upper left) and the main stages of the DDPM framework.
- Include a clear indication of the sampling process cycle.

- Label the 3D graph on the right hand side (e.g., “learned error distribution”).
- Reference these numbered steps in the text of Section 3.2 to link the figure with the equations in Section 3.1.

The revised figure is included in the revised manuscript.

We have also rephrased the corresponding paragraph:

“During training, the deterministic wind nowcast y is first used together with the analysis field to derive the forecast error x_0 , which represents the prediction target of the DDPM (Step ①). The error field x_0 is then gradually perturbed by Gaussian noise through the forward diffusion process, generating noisy error states x_t at different diffusion timesteps (Step ②). This process follows Eq. (1), where the original error distribution is progressively transformed into a Gaussian distribution.

For each diffusion timestep t , the noisy error x_t , the physical condition y , and the timestep embedding t are provided to the conditional denoising U-Net (Step ③). The denoiser estimates the noise component $\epsilon_\theta(x_t, y, t)$, which is used to reconstruct the forecast error during the reverse diffusion process. The network parameters are optimized by minimizing the difference between the true Gaussian noise and the predicted noise according to Eq. (4) (Step ④). Through this training procedure, the model learns the conditional distribution of forecast errors under different physical backgrounds y (Step ⑤).

During inference, the trained DDPM generates error realizations through the DDIM sampling strategy. Starting from random Gaussian noise $x_T \sim N(0, I)$, the model performs iterative reverse denoising from x_T to x_0 (Step ⑥). At each sampling step, the denoiser predicts the noise component conditioned on x_t , y , and t , and the DDIM update equation (Eq. (5)) is applied to obtain the previous diffusion state x_{t-1} . Repeating this reverse sampling process produces multiple possible realizations of the forecast error x_0 .

Finally, the generated error ensemble is combined with the deterministic wind nowcast to produce an ensemble wind forecast (Step ⑦):

$$y_{ens}^i = y + x_0^i, i = 1, \dots, N \quad (6)$$

where N denotes the number of generated ensemble members.”

The revision is in Line 139–158.

#Comments 6

150: „For verification, the SIVA forecast errors are designated as the reference for evaluating the DDPM, while the analysis fields serve as benchmark for the ensemble nowcasts.“

R: Did you use every grid-point of the SIVA domain and the analysis for verification? This is not clear, as in the caption of the Fig. 2 you denote observations as „Truth“, while in the caption of the Fig. 4. the analysis field is denoted „Truth“. It should be clarified in the Metrics description that where are you using point observations and where nowcasting software analysis for verification and why. It must be also taken into account that the analysis can already exhibit errors so it can be considered only as a near description of the real state of the atmosphere.

Reply 6:

Thanks for the comment.

We confirm that all verification uses every grid-point of the SIVA domain (e.g., error distributions, reliability diagrams, CRPS). In this study, the verification is based on the analysis field, which we treat as a gridded approximation of the true state (as noted in the response to

#Comment 1) and we do not use point-wise station observations for verification in the present work. We revised the sentence in Section 3.4 i.e.:

“In this study, we treat the gridded SIVA analysis field as a practical approximation of the true state. The forecast error of the deterministic nowcast is defined on the model grid as the difference between the SIVA forecast and the corresponding analysis field. For verification, we perform two separate evaluations. First, the errors generated by the DDPM are compared against the original SIVA forecast errors to assess how well the DDPM predicts these errors. Second, the ensemble nowcasts (obtained by adding the DDPM-generated errors to the SIVA forecast) are evaluated against the analysis field (as a surrogate for the true state) to assess the quality of the final probabilistic forecasts. All verification is performed on all grid points of the SIVA domain using standard probabilistic metrics. We note that the analysis is not a perfect truth; quantifying the impact of its residual errors on our results is left for future work, which will also include verification against independent station observations.” The revision is in Line 167–175.

We have replaced all instances of “truth” or “ground truth” with “analysis” throughout the manuscript (including figure captions) to avoid confusion between analysis and observations.

#Comments 7

155-160: „Given the negligible sensitivity of key verification metrics to ensemble size, an ensemble with 16 members was adopted to ensure statistical robustness while maintaining computational tractability.“

R: Can you support somehow your statement that the key verification metrics is not sensitive to ensemble size? Have you verified that? In the introduction, lines 35-45 you mention that „The ensemble often fails to fully characterize the true probability distribution“ or „However, a finite number of members cannot fully represent the true distribution, which inevitably leads to under-dispersion.“ This is seemingly in contradiction with your statement in the Evaluation Metrics part. How did you choose the 16 members then? Upon which criterion?

Reply 7:

Thanks for the comment. When designing the ensemble, we compared 8, 16, 32, and 50 members using CRPS, RMSE, and the RMSE/spread ratio. The figure 2 shows the results. The 8 members clearly underperforms the others. However, the differences among 16, 32, and 50 are minimal – CRPS differs by less than 0.01, and the RMSE/spread ratio differs by less than 0.02. Therefore, we chose 16 members as a balance between statistical robustness and computational cost (2 minutes per forecast cycle). Larger sizes (32–50) would take 4–8 minutes with virtually no improvement.

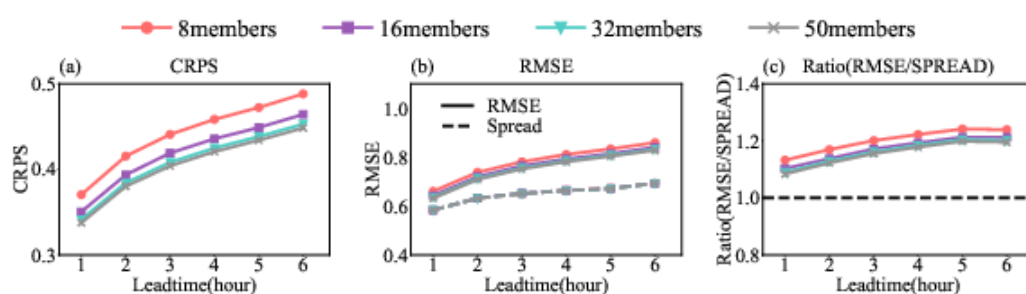


Fig.2. Verification metrics for ensemble nowcasts with different members over the test period. (a)

CRPS, (b) RMSE (solid lines) and ensemble spread (dashed lines). (c) Ratio of RMSE to SPREAD.

Regarding the apparent contradiction with the introduction: the description refers to classical ensemble methods where the true distribution is unknown and a finite sample inevitably leads to under-dispersion. But our model is trained to learn the target distribution of forecast errors. Once trained, generating an ensemble member is simply a random draw from that learned distribution. The distribution itself is already captured; increasing the number of draws beyond a certain point (here 16) brings negligible improvement. Thus, there is no contradiction. The introduction describes the general challenge, while our method addresses it by learning the distribution.

We have revised the text in Section 3.4 to clarify this reasoning and remove the phrase “negligible sensitivity” to avoid overstatement, i.e. “A comparison of ensemble sizes 8, 16, 32, and 50 shows that the 8-member ensemble underperforms the larger sizes, while the differences among 16, 32, and 50 members are minimal (e.g., CRPS differs by less than 0.01). Therefore, an ensemble size of 16 was chosen as a balance between statistical robustness and computational cost.” The revision is in Line 188–190.

#Comments 8

165-170: „The agreement of both the joint and marginal probability distributions with the benchmark (Fig. 2) demonstrates that the errors predicted by DDPM are physically consistent and statistically robust.“

R: Why do you think that if the statistical distribution of the evaluated model's forecasts and forecast errors matches the benchmark (which is considered to be physically consistent), then it must be physically consistent, too? Can you prove it? Or can you cite examples that if the physical consistency would not be fulfilled (e.g. using different approaches) then the forecast error distribution would be different?

Reply 8:

Thanks for the comment. We agree that matching the probability distributions alone does not prove physical consistency. Therefore, we have revised this sentence to state that the errors are “statistically consistent” rather than “physically consistent”, i.e. “The agreement of both the joint and marginal probability distributions with the benchmark (Fig. 3) demonstrates that the errors predicted by DDPM are statistically consistent.” The revision is in Line 199–200.

#Comments 9

205-210: First sentence of 4.2: „Evaluations of the generated errors reveal that DDPM captures the physical characteristics of forecast errors, thereby learning more than just their statistics.“

R: As for 165-170, I am not convinced that this was really shown as there was no comparison with other methods, where we would expect “only” learning statistics.

Reply 9:

Thanks for the comment. We agree that without comparison to other methods we cannot claim that the DDPM learns “physical characteristics” beyond statistics. Therefore, we have revised the sentence in Section 4.2 to simply state what the evaluations actually show, i.e., that the DDPM reproduces the statistics of the forecast errors. The revised sentence reads: “Evaluations of the generated errors reveal that the DDPM reproduces the statistics of the forecast errors.” The

revision is in Line 240.

#Comments 10

255-265 and a part of the Caption in Figure 7:

R: The discussion on the Brier Score and Brier Skill Score should be moved to 3.3 (Evaluation Metrics). Also the note on the use of climatology probability in BSS calculation. You could also mention, why the climatology probability was chosen to be a reference. Though a standard practice, it may have some implications for verification of rare events, especially when you used only the training dataset (~1,5 year).

Reply 10:

Thanks for the suggestions. We have made the following revisions:

1. Move the definitions of Brier Score and Brier Skill Score from the results section to Section 3.4 (Evaluation Metrics).

2. Add a note in Section 4.2 (Verification of Ensemble Nowcasts) i.e.: "It should be noted, however, that the climatological probability used as the reference for BSS is estimated from the training period (~1.5 years). This choice follows standard practice for probabilistic forecast verification (a no-skill benchmark). For rare events such as wind speeds exceeding 10.8 m s^{-1} , this short period yields a limited sample size, which may affect the robustness of the estimated climatology. This is a possible reason why the differences among the three noise schedules are less pronounced compared to the lower threshold. In future work, a longer climatological reference should be used to better evaluate probabilistic forecasts for rare events." The revision is in Line 308–314.

3. Remove the methodological sentence from the caption of this figure ("The reference value in calculating BSS is the climatology probability...") and move this information to Section 3.4, where the BSS definition is now given. The caption of this figure now only describes the results.

#Comments 11

285-295: Figure S2: You have quite a lot of description (one paragraph) concerning the output of Fig. S2 (Probability diagram of wind ensemble nowcasting in different lead time for three schedules with threshold 1 ms^{-1}). It would be fair to present it as a part of standard Figures of the manuscript, e.g. as Fig. 9a. The other diagram concerning the higher threshold of 10.8 ms^{-1} could be the Fig. 9b.

Reply 11:

Thanks for the suggestion. We have moved Figure S2 to the main text and merged it with the original Figure 9. The merged figure now is Figure 10, with panel (a) for the 1 m s^{-1} threshold and panel (b) for the 10.8 m s^{-1} threshold. The cross-references in the text have been corrected accordingly, i.e.: "Figure 10a shows the reliability diagrams for the 1 m s^{-1} wind speed, which lies at the centre of the observed probability distribution." and "The ensembles maintain high reliability at high wind events with threshold 10.8 m s^{-1} (Fig. 10b), exhibiting temporal evolution characteristics consistent with those observed for the 1 m s^{-1} threshold (Fig. 10a)." The revisions are in Line 332–337.

#Comments 12

295: You probably erroneously refer to Fig. S2, while the diagram in the current Fig.9 (for the threshold 10.8 m/s) is not referenced in the text at all. As mentioned above, consider moving Fig. S2 to Fig. 9a and denote current Figure 9 as Fig. 9b.

Reply 12:

Thanks for the hint. We have corrected the cross-references and now explicitly cite the high-threshold diagram (Fig. 9, now Fig. 10b after merging) in the text. Please see the response to [#comment 11](#) for details.

#Comments 13

315-325 and Figure 10: You mention Comparison with “truth” or “Ground Truth” although this is probably only the analysis of your deterministic nowcasting system and it may contain errors as I have mentioned earlier. It would be probably better to denote it as Analysis.

Reply 13:

Thanks for the suggestion. We have replaced “truth” and “ground truth” with “analysis” throughout the manuscript. The revision is traceable in the manuscript.

#Comments 14

Figure 10: A mistype occurred in the caption of the figure: Instead of “Ground Turth” there should be “Ground Truth” or even better “Analysis”.

Reply 14:

Thanks for pointing this out and sorry for the mistake. Following the suggestion, we have replaced all instances of “truth” or “ground truth” with “analysis” throughout the manuscript to avoid confusion. For the caption of this figure, we have corrected the typo as follows: “Fig. 11. Wind speed from 08:00 to 13:00 UTC on 10 June 2023. (a) Analysis, ...” The revision is in Line 381.

#Comments 15

Figure 10: “(c) Ensemble Mean forecast”

R: I do not understand why you did not show the forecasts of ensemble maxima, not even in the supplementary materials. Though, this is a parameter, which is often used in weather forecasting, especially for severe weather warnings. And you even note the “smoothing effect on the ensemble mean” in the text below the Figure (line 330). How could we know that the “strong wind band” considered as false alarm (319-320) and appearing in the deterministic nowcast (REF) on Fig. 10b would be not reproduced by some of the DDPM ensemble members? Only comparison with ensemble maximum of wind could exclude that. In the supplementary Fig. S3 one can see that although less expressed than in REF, but certain members (e.g. 9,10,12) show a signal for such wind band in that area.

Reply 15:

Thanks for the comment. We agree that stating the smoothing effect of the ensemble mean while calling the strong wind band in the deterministic forecast a false alarm is inconsistent

without checking the ensemble maximum. To resolve this, we have examined the ensemble maximum and added it to Figures 11 and 12 (new panels). The ensemble maximum in Figure 11 does show the strong wind band, but with lower probability. Therefore, we have revised the text accordingly i.e.: “The ensemble mean reduced the bias and captured the spatial structure more accurately, but the strong wind band was largely smoothed out (Fig. 11c). The ensemble maximum (Fig. 11d) and individual members (e.g., members 9, 10, 12 in Fig. S1) still show a weaker signal in that area. Thus, the band is not a complete false alarm; rather, the ensemble represents it with a non-zero but lower probability.” The revision is in Line 369–373.

#Comments 16

330-335 “The smoothing effect on the ensemble mean was also reduced, and the probabilistic forecast shows finer and more accurate details.”

R: It would be nice to mention an example. In addition, it would be perhaps noteworthy to highlight as an interesting feature that there is an area of reduced wind speed on the right edge of the domain (nearly in the middle, over the sea), visible in all outputs valid for the +1h time (Figure 11). Although it vanishes in both analysis and reference as the wind strengthens in time, it remains in the DDPM forecast until +6 hour, which suggests that the DDPM can exhibit certain “inertia” in some cases, even against its reference.

Reply 16:

Thanks for the suggestion. We have added a concrete example to support the statement. The example points to the wind speed centre over land near 33°N, 118°E, where the ensemble mean captures the location and intensity more accurately than the deterministic forecast. We also added a short discussion of the low-wind area over the sea near 33°N, 121.5°E. This feature fades in the analysis and reference after +1h but persists in the ensemble mean until +6h, while the ensemble maximum evolves more consistently with the analysis. We interpret this as a sign of under-dispersion in that region: the ensemble underestimates the probability of wind speed increase. This is consistent with the verifications (Fig. S3) that the model tends to be under-dispersive for extreme events. We have revised the text accordingly i.e. “Over the sea near 33°N, 121.5°E, a low-wind area is visible in all outputs at +1h. It fades in the analysis and deterministic reference after +1 hour as wind speed increases, but the ensemble mean retains it until +6 hour. The ensemble maximum, however, evolves more consistently with the analysis. This suggests that the ensemble members are not equally capturing the wind speed increase in this region, i.e., the ensemble is under-dispersive there. This behaviour is consistent with the under-dispersion noted for extreme events (Fig. S3).” The revision is in Line 390–396.

#Comments 17

350-355 “However, for extreme events, the model still suffers from the issue of excessive smoothness. This limitation may stem from the use of a basic diffusion model architecture.”

R: Is there not an additional problem that you used a relatively short (~1.5 year) training period? See my previous comment for the lines 85-90.

Reply 17:

Thanks for the hint. The relatively short training period (~1.5 years) is indeed an additional

factor contributing to the excessive smoothness for extreme events, as it likely contains too few samples of such events. We have revised the sentence in the conclusion to include this point. The sentence now reads: “This limitation may stem from the use of a basic diffusion model architecture as well as the relatively short training period (~1.5 years), which likely does not provide sufficient samples of extreme wind events.” The revision is in Line 436–441.

#Comments 18

370-375 “... while maintaining computational efficiency and physical consistency.”

R: Can you specify how much is DDPM computationally efficient? E.g. against a traditional nowcast (e.g. of your reference nowcast or some different nowcasts/ensembles).

Reply 18:

Thanks for the question. The deterministic SIVA nowcast takes about 2 minutes per forecast cycle, and our DDPM generates a 16-member ensemble in about 2 minutes on 2 NVIDIA 5090 GPUs. We have added computational efficiency information in the revised manuscript i.e.: “while maintaining computational efficiency: the deterministic reference nowcast takes about 2 minutes per forecast cycle, whereas the DDPM generates a 16-member ensemble in about 2 minutes on 2 NVIDIA 5090 GPUs.” The revision is in Line 463–464.

Referee comment #2

The manuscript introduces a paradigm for uncertainty quantification in short-term weather forecasting, specifically focusing on high-resolution, 10-meter wind speed nowcasting. Traditional uncertainty quantification relies heavily on dynamical ensemble prediction systems (EPS) that simulate multiple atmospheric trajectories by perturbing initial conditions or physical parameterizations. While effective, EPS is computationally intensive for real-time applications and frequently suffers from under-dispersion. To overcome these constraints, the authors propose bypassing the generation of multiple physical trajectories altogether. Instead, they leverage a generative artificial intelligence framework—specifically a Denoising Diffusion Probabilistic Model (DDPM)—to directly model and predict the full conditional distribution of forecast errors, using a deterministic physical forecast as the conditioning background.

Data preprocessing includes a logarithmic transformation of wind speed data, which successfully converts skewed wind speed errors into a near-Gaussian distribution and mathematically restricts the reconstructed wind speed fields to positive values. A core contribution of the paper is the empirical evaluation of three distinct noise scheduling frameworks (Linear, Cosine, and Sigmoid) within the diffusion process. The experimental results demonstrate that the choice of the noise scheduler is critical for structural and statistical fidelity. The Cosine schedule outperforms the alternatives across multiple verification metrics, achieving optimal error scores and well-calibrated histograms. Furthermore, spatial power spectrum analysis confirms that the samples generated via the Cosine schedule preserve the correct kinetic energy cascade and spatial structures of the wind fields.

The manuscript is well-structured, the physical treatment of data via logarithmic transformation is sound, and the comparison of noise schedules (Linear, Cosine, and Sigmoid) provides valuable empirical insights for the atmospheric modeling community. However, there are critical limitations regarding the temporal coverage of the dataset, the severe seasonal bias in the verification phase, the lack of traditional statistical baselines, mathematical inconsistencies in the description of the diffusion framework, and several data visualization/structural clarity issues that must be addressed before publication.

Reply:

We thank the reviewer for the thorough and constructive assessment of our work. We have carefully addressed all the concerns raised and provided clarifications where needed. The revisions are traceable in the manuscript.

Major Concerns:

#Comments 1

(Lines 85–87): The total historical record used spans from 1 October 2021 to 30 June 2023. This represents less than two full years of atmospheric data. Since deep generative models like DMs typically benefit from a substantial and diverse set of samples to properly map high-dimensional chaotic fields, a dataset shorter than two years may increase the risk of overfitting to the specific intra-annual anomalies of those two cycles, potentially limiting the model's capacity to generalize to long-term climate variability. It would be highly beneficial if the authors could provide a brief justification as to why a dataset shorter than 24 months is sufficient for high-resolution (1 km) spatial error generation, or alternatively, consider expanding the training set by incorporating additional years of archival data from the SIVA system.

Reply 1:

Thanks for the comments. We agree that a longer training period would generally be desirable, and we fully acknowledge the concern about the limited temporal coverage.

We would like to clarify an important distinction between our task and typical time-series forecasting with diffusion models. In many DM-based forecasting applications, the model is asked to predict a future state x_{T+t} given a past state x_T , meaning that the temporal evolution of the field must be learned from the data. In our case, the DDPM does not predict the future state directly; rather, it learns a conditional mapping from the deterministic forecast y (which is already valid at the target time) to the corresponding forecast error $e = x_{true} - y$. The input and output are contemporaneous—both refer to the same valid time. The temporal dimension enters only through the training samples, not through the prediction task itself. This conditional formulation alleviates the need to learn temporal dynamics, and consequently reduces the demand for a long continuous time series.

To provide some empirical evidence that the model can generate spatial errors sufficiently, we compared the daily-mean wind speed from the DDPM ensemble mean, the deterministic SIVA forecast, and the analysis field for the test month (January and June 2023). As shown in Figure 1 and 2, the ensemble mean consistently lies closer to the analysis than the deterministic forecast does, indicating that the model learns useful error structures even with the current dataset.

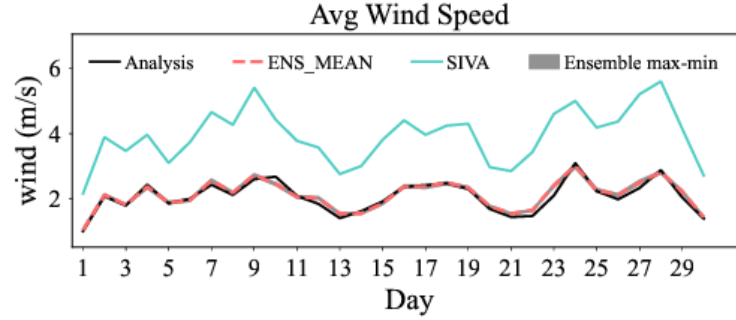


Fig.1. Daily mean wind speed over the test period (June 2023) at lead time 6h, comparing the analysis field (black), DDPM ensemble mean with cosine noising schedule (dashed orange), SIVA nowcast (light blue), and the ensemble maximum and minimum (shadow).

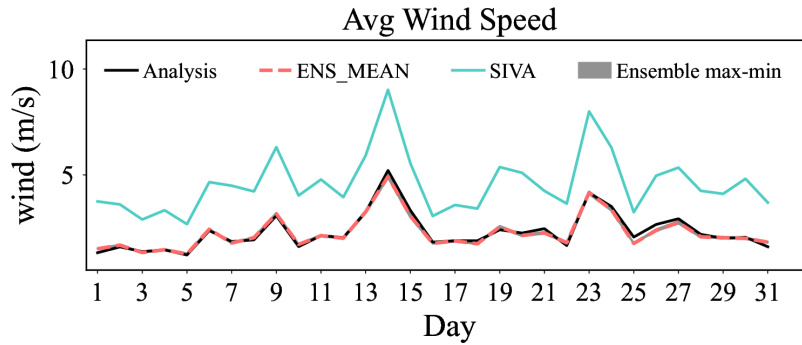


Fig.2. The same as Fig1 but for January 2023.

We acknowledge that a longer training period would be beneficial, especially for rare extreme events. This limitation has been noted in the Conclusions i.e.: “This limitation may stem from the use of a basic diffusion model architecture as well as the relatively short training period (~1.5

years), which likely does not provide sufficient samples of extreme wind events.” The revision is in Line 432–433.

#Comments 2

(Lines 88–90 and Line 167): The authors state that the dataset was split chronologically, leaving the "remainder" for testing (which corresponds strictly to June 2023, as confirmed in Line 167). As defined in Lines 85–86, the spatial domain covers East China (115.35°E – 122.33°E , 29.88°N – 35.81°N). In this region, June is climatologically dominated by a meteorological regime characterized by stationary front dynamics. Evaluating the DDPM's performance exclusively during this single month could introduce a seasonal bias. Unless a specific seasonal behavior is being targeted (which should then be explicitly mentioned in the manuscript), it is recommended to evaluate the model across a broader range of seasonal conditions. Given the relatively small size of the dataset, a comprehensive evaluation might be challenging; however, adopting a validation strategy that ensures cross-seasonal representation would greatly help to demonstrate the model's generalizability across different atmospheric regimes.

Reply 2:

Thanks for the comment. We have retrained the model using a training period that excludes the test months: October 2021 – December 2022 and February 2023 – April 2023. We evaluated the Cosine schedule (the best-performing scheme) on both January 2023 (winter) and June 2023 (summer). The results show that the model performs consistently across the two seasons, with no significant degradation in winter performance. The seasonal evaluation has been added to Section 4.4. in Line 407–420.

#Comments 3

(Lines 68–70 and Lines 159–161): In Lines 68–70, the authors state that, unlike conventional statistical post-processing methods, their proposed model explicitly learns the full conditional distribution of errors. Later, in Lines 159–161, they explain the lack of a reference ensemble by noting that the operational SIVA system is purely deterministic. While this clarifies the absence of a dynamical ensemble comparison, it would be highly valuable to include a reference baseline for probabilistic verification. To more clearly demonstrate the added value and skill of the conditional estimation provided by the DDPM, the model could be benchmarked against standard reference baselines. Even a simple climatological distribution or a persistence-based error model would serve as an insightful benchmark to quantify the actual statistical improvement and contextualize the computational investment required by the diffusion architecture.

Reply 3:

Thanks for the suggestion. We agree that a reference baseline would help to demonstrate the added value of the DDPM. We implemented a climatological baseline, where ensemble members are generated by randomly sampling from the historical distribution of SIVA forecast errors (pooled from the training set). The performances of this climatological baseline have been compared to the DDPM using the same verification metrics (e.g., CRPS, Brier score, reliability diagrams, et al.). The results are presented and discussed in Section 4. A brief description of the baseline method is added in Section 3.3 i.e. “To provide a reference for evaluating the

DDPM-based ensemble, a climatological baseline is constructed. For each grid point, the mean and variance of the forecast errors are estimated from the training set. Since the error distributions at individual grid points are approximately Gaussian, the baseline generates a 16-member ensemble for each lead time by randomly sampling from the corresponding Gaussian distribution at each grid point and adding the sampled errors to the deterministic SIVA nowcast. This baseline does not use any conditional information beyond the climatological error statistics at each grid point.” The revision is in Line 160–165.

#Comments 4

(Lines 105–106): In Lines 105–106, the text states that $\mu_\theta(x_t, t)$ and σ^2 are the mean and variance of the distribution at the previous step. While the explicit formulation to obtain μ_θ is provided in Equation 3, it would be helpful if the authors could explicitly define how σ^2 is handled or computed.

Reply 4:

Thanks for the comment. In our implementation, we follow Ho et al. (2020) and set $\sigma^2 = \beta_t$, where β_t is the predefined noise variance at step t (the same as in the forward process). We have added a brief clarification after Equation (2) i.e.: “we set $\sigma^2 = \beta_t$, i.e., the same variance as in the forward process at step t .” The revision is in Line 111.

#Comments 5

(Line 109 and Equation 4): In Line 109, the network is introduced as $\epsilon_\theta(x_t, t)$, yet in Equation 4, it appears as $\epsilon_\theta(x_t, y, t)$. To avoid ambiguity, the authors should clarify whether the conditional variable y should be consistently included in the notation throughout the text.

Reply 5:

Thanks for pointing it out. The conditional variable y (the deterministic nowcast) is used throughout the diffusion process. We have consistently revised $\epsilon_\theta(x_t, y, t)$ in both the text and Equation 4 to avoid ambiguity, i.e.: “where $\epsilon_\theta(x_t, y, t)$ is a function approximator to predict the noise ϵ that conditioned on the deterministic nowcast y and the current step t .” The revision is in Line 114–115.

#Comments 6

(Figure 1 / Section 2.3): The flowchart in Figure 1, which outlines the overall framework of the model, can be somewhat difficult to follow, even when cross-referenced with the text explanation in Section 2.3. To further enhance its clarity and make it more intuitive for the reader, it is suggested to refine the representation of the data flow. Specifically, explicitly linking the visual blocks to the mathematical variables used in the text—such as the physical background (y), the forward/reverse processes (x_t), and the final output—would significantly improve the diagram's transparency. Explicitly referencing the corresponding equation numbers (Equations 1–4) within the flowchart components would also greatly assist the community in understanding and reproducing the exact methodology.

Reply 6:

Thanks for the suggestion. We have redesigned the flowchart to make it more self-explanatory.

Specifically:

- Add numbered steps (①, ②, ...) to indicate the start of training and the main stages.
- Explicitly link the visual blocks to mathematical variables.
- Reference the corresponding equation numbers within the flowchart.

The revised figure is included in the manuscript.

We have also rephrased the corresponding paragraph in Line 139–158:

“During training, the deterministic wind nowcast y is first used together with the analysis field to derive the forecast error x_0 , which represents the prediction target of the DDPM (Step ①). The error field x_0 is then gradually perturbed by Gaussian noise through the forward diffusion process, generating noisy error states x_t at different diffusion timesteps (Step ②). This process follows Eq. (1), where the original error distribution is progressively transformed into a Gaussian distribution.

For each diffusion timestep t , the noisy error x_t , the physical condition y , and the timestep embedding t are provided to the conditional denoising U-Net (Step ③). The denoiser estimates the noise component $\epsilon_\theta(x_t, y, t)$, which is used to reconstruct the forecast error during the reverse diffusion process. The network parameters are optimized by minimizing the difference between the true Gaussian noise and the predicted noise according to Eq. (4) (Step ④). Through this training procedure, the model learns the conditional distribution of forecast errors under different physical backgrounds y (Step ⑤).

During inference, the trained DDPM generates error realizations through the DDIM sampling strategy. Starting from random Gaussian noise $x_T \sim N(0, I)$, the model performs iterative reverse denoising from x_T to x_0 (Step ⑥). At each sampling step, the denoiser predicts the noise component conditioned on x_t , y , and t , and the DDIM update equation (Eq. (5)) is applied to obtain the previous diffusion state x_{t-1} . Repeating this reverse sampling process produces multiple possible realizations of the forecast error x_0 .

Finally, the generated error ensemble is combined with the deterministic wind nowcast to produce an ensemble wind forecast (Step ⑦):

$$y_{ens}^i = y + x_0^i, i = 1, \dots, N \quad (6)$$

where N denotes the number of generated ensemble members.”

Minor Concerns:

#Comments 7

(Figure 6 / Section 4.1): In Section 4.1, specific conclusions are drawn based on the spatial patterns shown in Figure 6. However, the current color palette utilizes very soft gradients and smooth transitions, which can make it challenging for the reader to visually discern the differences highlighted in the text. It is suggested that the authors consider changing the color scale of Figure 6 to a high-contrast or perceptually uniform divergent palette to visually enhance and further substantiate the claims made in the manuscript.

Reply 7:

Thanks for the suggestion. We have changed the color palette of the mentioned figure to a perceptually uniform diverging scheme to better highlight the spatial differences discussed in the text. The revised figure is included in the manuscript and now the number is Figure 7.

#Comments 8

(Supplementary Material / Text Structure): The manuscript contains a prolonged text explanation discussing a specific figure that has been placed entirely in the Supplementary Material. When a figure requires such an extensive textual breakdown and is central to supporting the core arguments, forcing the reader to search for it in a separate document can disrupt the flow of reading. To make the paper more self-contained and streamline the reading process, it is recommended to move this specific figure from the supplementary material into the main manuscript, embedding it close to its corresponding text discussion.

Reply 8:

Thanks for the suggestion. We agree that a figure requiring detailed discussion should not be placed in the supplementary material. We have moved this figure to the main text and updated the cross-references accordingly, i.e.: “Figure 10a shows the reliability diagrams for the 1 m s⁻¹ wind speed, which lies at the centre of the observed probability distribution.” and “The ensembles maintain high reliability at high wind events with threshold 10.8 m s⁻¹ (Fig. 10b), exhibiting temporal evolution characteristics consistent with those observed for the 1 m s⁻¹ threshold (Fig. 10a).” The revisions are in Line 332–340.