

Methodological Factors Affecting the Generalizability of Deep Learning-Based Digital Elevation Model Super-Resolution

Zahra Anvarian¹, Farhad Maleki¹

¹Department of Computer Science at University of Calgary

*Correspondence: mahsa.anvarian@ucalgary.ca; farhad.maleki1@ucalgary.ca

1 Abstract

2 High-resolution digital elevation models (DEMs)
3 are important for terrain-dependent analysis in
4 precision agriculture and remote sensing. Deep
5 learning-based DEM super-resolution (SR) has
6 shown promising results for generating higher-
7 resolution DEMs from lower-resolution inputs;
8 however, reported performance can be strongly
9 affected by how training and test data are se-
10 lected. This paper presents a methodological
11 evaluation of DEM SR protocols, focusing on
12 spatial leakage, test-region representativeness,
13 and dataset-level terrain diversity, three key
14 factors that can substantially limit the gener-
15 alizability of DEM SR models. Using 30 m
16 SRTM DEM data, we generate 120 m-to-30
17 m SR pairs and evaluate the TfaSR (Terrain
18 feature-aware Super-Resolution) model under
19 controlled experimental conditions to isolate the
20 effect of each factor. First, randomly splitting
21 tile-based data into training and test sets in-
22 troduces spatial leakage that inflates apparent
23 performance: the model achieved a DEM RMSE
24 of 4.52 m on the internal random test set, while
25 evaluation on a terrain-matched but geographi-
26 cally independent test region raised this to 5.02
27 m, revealing that random-split evaluation over-
28 estimates true generalization. Second, a low-
29 relief test region produced substantially lower
30 DEM RMSE than a terrain-representative test
31 region (1.68 vs. 5.02 m), a gap that reflects
32 the simpler terrain structure of flat areas rather
33 than superior model generalization. Third, on a
34 shared terrain-diverse test set, TfaSR trained on
35 a terrain-diverse dataset achieved lower DEM
36 RMSE than the official pretrained TfaSR, whose
37 training data is geographically distributed but

38 low-relief-dominated (5.50 vs. 5.69 m), indicat-
39 ing that terrain diversity in the training data,
40 rather than broad geographic coverage alone,
41 drives generalization to varied terrain. These
42 results show that random tile-based splitting,
43 unrepresentative test-region selection, and lim-
44 ited training-data terrain diversity can lead to
45 misleading conclusions about DEM SR gener-
46 alization. Credible evaluation should explicitly
47 report spatial partitioning strategies, train-test
48 terrain-feature similarity, and dataset-level ter-
49 rain composition.

50 **Keywords:** Digital elevation models (DEMs),
51 Super-resolution, Geospatial deep learning models,
52 Spatial leakage, Terrain representativeness

53 1 Introduction

54 High-resolution digital elevation models (DEMs)
55 are essential geospatial data products for terrain-
56 dependent analysis in precision agriculture and
57 remote sensing, supporting tasks such as hydro-
58 logical modelling, irrigation and drainage plan-
59 ning, erosion risk assessment, and terrain char-
60 acterization¹⁻³. Many of these tasks depend on
61 fine-scale elevation variation and local terrain
62 structures that coarse-resolution DEMs may not
63 adequately capture, yet global DEM products
64 often lack the required spatial resolution^{4,5}. Be-
65 cause acquiring high-resolution DEMs through
66 LiDAR, photogrammetry, or field surveying is of-
67 ten expensive or impractical, there is increasing
68 interest in DEM super-resolution (SR), which
69 reconstructs a higher-resolution DEM from a
70 lower-resolution input^{6,7}.

71 Recent deep learning-based DEM SR meth-
72 ods have reported promising results by adapt-
73 ing image super-resolution architectures, in-

cluding convolutional neural networks, residual networks, generative adversarial networks, and transformer-based models^{8–14}. In these approaches, DEMs are commonly treated as single-channel raster data, and models are trained to predict high-resolution elevation grids from degraded low-resolution inputs. Performance is typically evaluated using reconstruction metrics such as root mean square error (RMSE), mean absolute error (MAE), or peak signal-to-noise ratio (PSNR). While these metrics are useful, they do not by themselves establish whether a model generalizes reliably to geographically independent terrain. For geospatial raster data, reported performance depends not only on model architecture and loss functions, but also on how training and test samples are selected.

A major concern is spatial leakage caused by random tile-based train/test splitting. Because DEM tiles are sampled from continuous, spatially autocorrelated terrain surfaces, randomly selected test tiles may lie close to training tiles and share highly similar topographic structure^{15–17}. Test results obtained under such splits may therefore overstate how well a model generalizes to independent geographic regions, even though tile-based datasets and random partitioning remain common in DEM SR studies^{12–14}.

A second concern is the representativeness of the selected test region. Even a geographically separated test area can be artificially easy if it is dominated by low-relief terrain, or artificially difficult if it is substantially more rugged than the training data. Reliable evaluation therefore requires characterizing test regions in terms of terrain-feature distributions, such as elevation variation, slope, Terrain Ruggedness Index (TRI), and Topographic Position Index (TPI), rather than relying on geographic separation alone.

A related concern is the terrain diversity of the dataset as a whole. A dataset spanning multiple geographic areas may still represent a narrow terrain domain if most samples come from flat or low-relief landscapes, in which case strong results on a random split only support generalization within that restricted domain. Eval-

uating dataset-level terrain diversity therefore requires examining the distribution of terrain features across samples, rather than relying on geographic coverage as a proxy for topographic breadth.

This paper presents a methodological evaluation of deep learning-based DEM SR, focusing on evaluation validity rather than architectural novelty. While each of these concerns has been recognized in the broader geospatial machine learning literature, their quantified impact on DEM SR evaluation has not been systematically examined under controlled experimental settings. Using 30 m SRTM (Shuttle Radar Topography Mission) DEM data from Alberta and British Columbia, Canada, we construct 120 m-to-30 m SR pairs and evaluate the TfaSR model¹². To isolate the effect of evaluation protocol, the model and training configuration are kept fixed across all experiments.

The main contributions of this paper are as follows:

1. We quantify the effect of random tile-based splitting in DEM SR by comparing internal random testing with geographically independent region-based testing under a fixed model configuration.
2. We show that region-based testing alone is insufficient unless the test region is terrain-representative of the training region, and we evaluate representativeness using terrain-feature distributions and similarity measures.
3. We demonstrate with controlled experiments that geographic coverage alone does not ensure terrain diversity: a training dataset collected from multiple regions may still be dominated by low-relief terrain, limiting generalization to more varied topographic conditions.

We argue that credible DEM SR evaluation should move beyond random tile splits and aggregate reconstruction metrics, explicitly reporting spatial partitioning strategies, train-test terrain-feature similarity, and dataset-level terrain di-

versity so that generalization and real-world applicability claims are better supported.

2 Related Work

Deep learning-based DEM SR has advanced by adapting architectures originally developed for image SR, including convolutional networks, residual networks, generative adversarial networks, and transformer-based models^{8–11}. DEM-specific methods have further introduced terrain-aware components to better preserve topographic structure. D-SRGAN and D-SRCAGAN applied adversarial learning and channel attention, showing that learned models can outperform traditional interpolation by recovering nonlinear relationships between low- and high-resolution elevation grids^{6,18}. TfaSR incorporated deformable convolution and terrain-related loss terms to improve the reconstruction of local terrain features such as ridges and valleys¹². GISR introduced global-information constraints to improve consistency between local and broader elevation patterns¹³, and transformer-based approaches such as DSRT and TTSR aimed to capture local–global dependencies and multiscale terrain heterogeneity^{14,19}. These studies demonstrate important progress in DEM SR model design; however, most rely on random tile-based partitioning within a fixed geographic region and evaluate results using aggregate reconstruction metrics, without examining whether reported improvements support reliable generalization to independent terrain.

A central issue in geospatial model evaluation is spatial dependence between samples. Nearby locations often share similar terrain characteristics due to spatial autocorrelation¹⁵, and prior work in remote sensing has shown that random sample partitioning can overestimate predictive performance when spatial dependence is ignored^{16,17,20}. This concern applies directly to DEM SR, where tiled datasets are generated from continuous terrain surfaces and random tile-based splitting is the prevailing evaluation practice^{12–14}. DEM SR has largely inherited image SR’s evaluation conventions, in which random splits are standard and sample independence is assumed — an assumption that does not hold for spatially

autocorrelated terrain data.

Beyond spatial separation, the terrain characteristics of the test data also affect the interpretation of DEM SR results. Terrain analysis literature provides established descriptors such as elevation variability, slope, TRI, and TPI for characterizing terrain complexity^{1–3,21,22}; however, DEM SR studies typically report performance without explicitly characterizing training terrain similarity or dataset-level terrain diversity^{6,12–14,18}. Together, these gaps suggest that DEM SR generalization remains understudied as an evaluation problem, independently of the architectural advances being reported.

3 Methodological Evaluation Framework

The evaluation of DEM SR differs from standard image SR because DEM samples are spatially referenced and represent continuous physical terrain surfaces. As a result, the validity of reported performance depends not only on the model and reconstruction metrics, but also on the spatial partitioning strategy, the terrain characteristics of the test region, and the terrain diversity of the dataset, which the following subsections examine in turn.

3.1 Random Tile-Based Splitting and Spatial Leakage

A common evaluation protocol in DEM SR divides tiled DEM samples into training, validation, and test subsets by random sampling, assuming that individual tiles are sufficiently independent. However, DEM tiles are extracted from continuous terrain surfaces, where nearby samples may share similar elevation ranges, slope patterns, ridges, valleys, drainage structures, and surface roughness; randomly selected test tiles may therefore be geographically close to training tiles and contain highly similar terrain patterns.

Figure 1 illustrates this issue using the DEM SR dataset released with TfaSR¹², which is publicly available through Figshare (DOI: <https://doi.org/10.6084/m9.figshare.19225374>). Although the dataset includes multiple geographic regions, random partitioning is applied at the tile level within the sampled areas. As shown

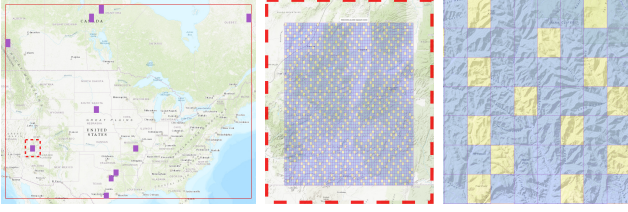


Figure 1: Random tile-based splitting in the publicly released TfaSR DEM SR dataset¹². Yellow tiles are test samples; blue tiles are training samples. The zoomed panel shows test tiles directly adjacent to training tiles.

We define test-region representativeness as the degree to which a geographically independent test region reflects the terrain-feature distribution of the training region for the intended evaluation objective. This does not require the test region to be identical to the training region, but it should be comparable enough that the evaluation is not artificially easy, artificially difficult, or unrelated to the claimed application domain. We assess representativeness using terrain-feature distributions based on elevation variation, slope, TRI, and TPI. Our second experiment compares a low-relief external test region with a terrain-representative external test region to quantify how test-region selection affects reported DEM SR accuracy.

3.3 Dataset-Level Terrain Diversity

Train-test similarity should also be interpreted in relation to the terrain diversity of the dataset as a whole. A dataset may include multiple geographic areas but still represent a narrow terrain domain if most samples are concentrated in flat or low-relief landscapes. In such cases, a random split may produce training and test sets with similar feature distributions, but the resulting evaluation only supports generalization within that limited terrain domain.

We define dataset-level terrain diversity as the range of terrain conditions represented across the full evaluation dataset. This factor is important because geographic coverage does not necessarily imply topographic diversity. To support broad generalization claims, a DEM SR dataset should contain a sufficient range of terrain complexity, or the scope of the claims should be restricted to the terrain domain represented by the data. Our third experiment examines this factor by comparing a model trained on a geographically distributed dataset with one trained on a terrain-diverse dataset of equal size, evaluated on a shared, geographically independent test set, to test whether geographic coverage alone is sufficient for broad generalization.

4 Terrain-Feature-Based Analysis

To evaluate whether DEM SR test results support meaningful generalization claims, we char-

acterize the terrain conditions of the training and test data at the tile level, using descriptors computed from the reference high-resolution DEM. These descriptors serve both to assess test-region representativeness and to gauge dataset-level terrain diversity.

4.1 Terrain Descriptors

Each DEM tile is represented using a set of terrain descriptors that summarize both elevation variation and local terrain structure. Elevation-based descriptors include mean elevation, standard deviation, elevation range, interquartile range, 5th–95th percentile range, and median absolute deviation. These measures describe the vertical variation within each tile and help distinguish low-relief tiles from tiles with stronger elevation variability. However, elevation statistics alone do not fully characterize terrain complexity because two tiles may have similar elevation ranges but different local surface forms.

To capture local terrain structure, we compute slope, TRI, and TPI^{2,21,22}. Slope describes the local rate of elevation change, TRI measures local ruggedness by comparing a cell with its surrounding neighbourhood, and TPI measures whether a cell is locally higher or lower than its neighbouring cells. For a cell with elevation z_c and eight neighbouring elevations z_i , TRI and TPI are defined as:

$$TRI = \sqrt{\sum_{i=1}^8 (z_c - z_i)^2}, \quad TPI = z_c - \frac{1}{8} \sum_{i=1}^8 z_i. \quad (1)$$

For each tile, the derived slope, TRI, and TPI layers are summarized using mean, standard deviation, and maximum values. Together, these per-tile elevation and slope/TRI/TPI summaries form the terrain-feature set on which all similarity measures reported in this paper are computed.

4.2 Train-Test Similarity Measures

After computing terrain descriptors for all tiles, we compare the feature distribution of each candidate set against a reference set chosen per experiment (the internal random test set in Experiment 1 and the training set in Experiment 2). A standard scaler is fitted on the reference set

only and applied to the compared set, so comparisons are made in a common feature space without using the compared set’s information during normalization.

We use complementary similarity measures rather than relying on a single criterion. The mean standardized feature shift measures the difference between the reference and compared feature means after scaling. The Wasserstein distance and Kolmogorov–Smirnov (KS) statistic compare the two distributions feature by feature, where lower values indicate stronger distributional similarity^{23,24}.

Finally, we use classifier-based separability as a multivariate similarity measure²⁵. A logistic regression classifier is trained to distinguish training tiles from test tiles using only terrain descriptors. The resulting area under the receiver operating characteristic curve (AUC) is converted into a distance score:

$$d_{AUC} = 2|AUC - 0.5|. \quad (2)$$

A value near zero indicates that the classifier cannot reliably distinguish training and test tiles, suggesting high terrain-feature similarity. A value near one indicates that the two sets are easily separable and therefore less similar in the selected terrain-feature space.

These measures are interpreted jointly: a representative test region should show low mean feature shift, low distributional distance, and low classifier separability. The same descriptors also support dataset-level analysis of whether the full dataset spans a broad range of terrain conditions or is concentrated within a narrow domain.

5 Experiments and Results

5.1 Experimental Setup

All model-based experiments use 30 m SRTM DEM data and follow a 120 m-to-30 m DEM SR setting. High-resolution (HR) DEM tiles are extracted as non-overlapping 64×64 patches. Low-resolution (LR) inputs are generated by down-sampling the HR tiles by a factor of four using nearest-neighbour interpolation. This degradation process ensures spatial alignment between LR inputs and HR reference DEMs.

We use TfaSR as the reference DEM SR model because it is a terrain-aware deep learning architecture designed specifically for DEM SR. To isolate the effect of evaluation protocol and dataset composition, all experiments use the same model architecture, loss function, normalization strategy, and training hyperparameters. The model is trained using the combined content mean squared error (MSE) and slope MSE loss, Adam optimizer, learning rate of 10^{-4} , and batch size of 32. In all experiments, models are trained for 100 epochs, and the final checkpoint is used for evaluation to avoid test-set-based model selection.

Performance is reported using DEM RMSE, slope RMSE, and aspect RMSE. DEM RMSE measures elevation reconstruction error directly, while slope and aspect RMSE provide additional information about errors in terrain derivatives. Experiment 3 additionally reports TRI and TPI RMSE, which measure reconstruction error in the TRI and TPI derivative layers. Aspect RMSE is interpreted with caution in low-relief regions because aspect can become unstable when local slope values are close to zero.

5.2 Experiment 1: Random Tile-Based Testing vs. Region-Based Testing

The source study area is located in Alberta and British Columbia, Canada, and is defined by the bounding box $[-118.6414, 49.5252, -115.0269, 51.9172]$.¹ From this region, 26,954 HR DEM tiles were generated. These tiles were randomly divided into 21,563 training tiles and 5,391 internal test tiles using an 80/20 random split.

As a more appropriate evaluation for this task, we also extracted a separate test region defined by the bounding box $[-120.1849, 52.1301, -117.7734, 53.0521]$. This external region was not used during training and is geographically separated from the source training area. It was subsampled to match the internal random test set in both

sample count (5,391 tiles) and terrain-feature distribution (verified in Figure 2 and Table 1), so that performance differences reflect spatial separation rather than terrain difficulty. The same trained model was evaluated on both the internal random test set and the external region-based test set.

Figure 2 compares the terrain-feature distributions of the training set, the internal random test set, and the external region-based test set using elevation range, mean TRI, mean TPI, and mean slope. The internal random test set closely follows the training distribution, as expected from random sampling within the same source region. Importantly, the external region-based test set also shows substantial overlap with the training distribution across these descriptors, while remaining geographically independent. This makes it suitable for evaluating whether performance changes when the test data are spatially separated rather than randomly sampled from the same continuous region.

Table 1 complements this visual comparison with quantitative terrain-feature similarity measures computed relative to the internal random test set, for the external region-based test set, with the training set shown as context. These measures determine whether the external region provides terrain comparable to the internal random test set while remaining geographically independent.

Table 2 summarizes the DEM SR performance under the two evaluation settings. The internal random test set produces lower DEM RMSE than the external region-based test set. Since the training configuration and model are fixed, this difference suggests that random tile-based evaluation can provide a more optimistic estimate of performance than testing on a geographically independent region.

Because Table 1 confirms that the external region is terrain-matched to the internal random test set, this performance gap is associated mainly with spatial independence rather than terrain difference.

¹Bounding boxes are reported in decimal degrees as [west, south, east, north], that is, minimum longitude, minimum latitude, maximum longitude, and maximum latitude.

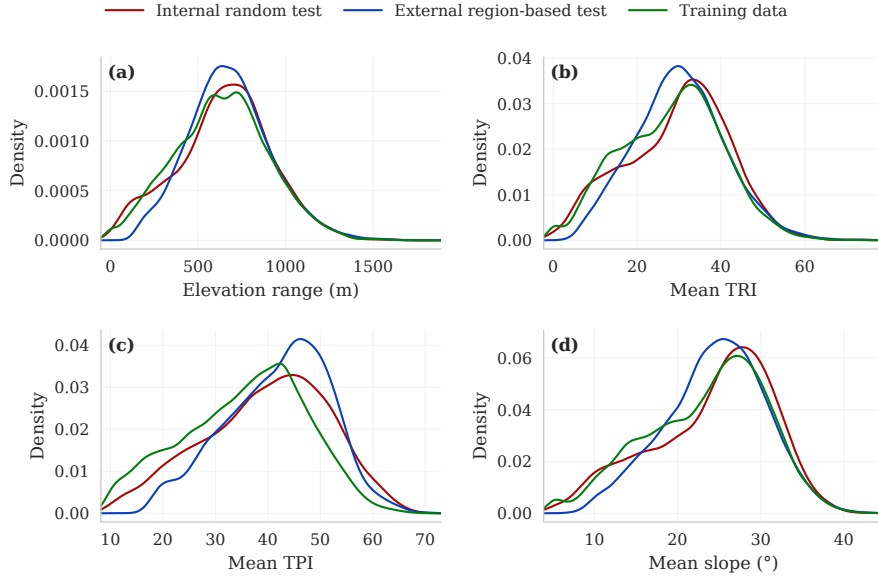


Figure 2: Terrain-feature distributions for Experiment 1: training set, internal random test set, and external region-based test set.

Table 1: Terrain-feature similarity in Experiment 1, relative to the internal random test set (training set shown as context). Lower values indicate stronger similarity.

Compared set	Mean shift	Wasserstein	KS	d_{AUC}
External region-based test	0.088	0.117	0.063	0.218
Training data (context)	0.221	0.186	0.092	0.255

5.3 Experiment 2: Low-Relief vs. Terrain-Representative Test Region

We compare two geographically independent test regions. The first is a low-relief test region defined by the bounding box $[-118.0041, 53.3780, -115.9058, 54.4125]$, which contains limited elevation variation and smoother terrain. The second is the external region-based test set from Experiment 1, which is selected to better match the terrain-feature distribution of the training region.

Both test regions are evaluated using the same trained model from Experiment 1, isolating the effect of test-region terrain complexity. Figure 3 compares the terrain-feature distributions of the training set and the two candidate test regions using elevation range, mean TRI, mean TPI, and mean slope. The low-relief test region is concentrated at substantially lower values across all four descriptors, whereas the terrain-representative test region closely overlaps the training distribution. Table 3 reports the corre-

sponding terrain-feature similarity between the training set and each candidate test region.

Table 4 shows that the low-relief test region produces substantially lower DEM RMSE and slope RMSE than the terrain-representative test region. This result indicates that test-region selection can strongly affect reported accuracy. A low-relief test region may produce favourable error values because the reconstruction task is less complex, not necessarily because the model generalizes robustly across more complex terrain.

The high aspect RMSE in the low-relief region should be interpreted carefully. In flat or nearly flat areas, aspect is often unstable because small elevation differences can cause large angular changes in the estimated aspect direction. Therefore, the low DEM RMSE and slope RMSE are more informative for assessing the relative ease of the low-relief test region.

Table 2: DEM SR performance: internal random vs. external region-based testing.

Evaluation setting	DEM RMSE (m)	Slope RMSE ($^{\circ}$)	Aspect RMSE ($^{\circ}$)
Internal random test	4.52	2.99	50.51
External region-based test	5.02	3.29	50.34

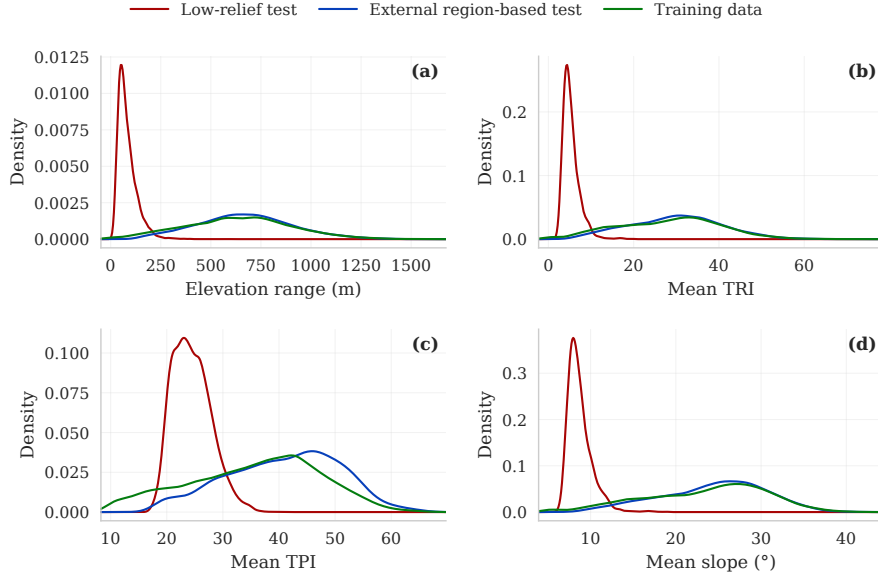


Figure 3: Terrain-feature distributions for Experiment 2: training set, low-relief test region, and terrain-representative test region.

5.4 Experiment 3: Effect of Dataset-Level Terrain Diversity on Generalization

We compare the official pretrained TfaSR model, trained on a geographically distributed but low-relief-dominated dataset, with a TfaSR model trained on a terrain-diverse dataset of equal size. The terrain-diverse training dataset spans four geographically separated regions covering mountainous, alpine, and moderate hilly terrain, and the test set spans three geographically independent regions that do not overlap with the training regions or with the original TfaSR dataset.² To make the comparison consistent in terms of dataset size, we selected 25,088 training patches and 6,272 test patches, matching the split size of the TfaSR dataset. The terrain-diverse TfaSR

²Training bounding boxes: $[-118.6414, 49.5252, -115.0269, 51.9172]$, $[85.4901, 27.3000, 87.4786, 28.4304]$, $[87.6544, 27.4400, 88.8739, 28.1117]$, and $[-119.4653, 47.7098, -117.2481, 48.9874]$. Test bounding boxes: $[-120.1849, 52.1301, -118.20, 53.0521]$, $[94.2435, 30.3515, 95.2213, 30.7586]$, and $[85.0342, 49.1155, 87.9565, 50.9030]$.

model was trained using the same architecture and training configuration as the original TfaSR model, including the same degradation process, tile size, loss function, optimizer, learning rate, batch size, normalization strategy, and number of epochs; only the training dataset was changed. The official pretrained TfaSR model was then evaluated directly on the same terrain-diverse test set. Because the only difference between the two models is the terrain composition of their training data, the comparison isolates the effect of training-data terrain diversity.

Figure 4 compares the terrain-feature distributions of the original TfaSR dataset and the terrain-diverse dataset. The TfaSR dataset shows a strong concentration in low-relief terrain, despite being collected from multiple geographic areas. In contrast, the terrain-diverse dataset covers a wider range of elevation variation, slope, TRI, and TPI values, indicating broader terrain diversity. The terrain-diverse training and test distributions also show substantial overlap, suggesting that the test set is representative of the terrain conditions in the training data while re-

Table 3: Terrain-feature similarity between the training set and each candidate test region in Experiment 2. Lower values indicate stronger similarity.

Test region	Mean shift	Wasserstein	KS	d_{AUC}
Low-relief test region	1.514	1.461	0.769	0.995
Terrain-representative test region	0.266	0.232	0.107	0.375

Table 4: DEM SR performance: low-relief vs. terrain-representative test region.

Test region	DEM RMSE (m)	Slope RMSE ($^{\circ}$)	Aspect RMSE ($^{\circ}$)
Low-relief test region	1.68	1.44	96.20
Terrain-representative test region	5.02	3.29	50.34

1 maining geographically independent.
2 Table 5 reports the performance of both models
3 on the shared terrain-diverse test set. The model
4 trained on the terrain-diverse dataset achieves
5 lower error across all reported metrics. In terms
6 of DEM RMSE, it improves from 5.69 m to 5.50
7 m compared with the official pretrained TfaSR
8 model. Similar improvements are observed for
9 TRI, TPI, slope, and aspect RMSE. Although
10 the improvement is moderate in absolute terms,
11 it is statistically significant and consistent: a
12 one-sided Wilcoxon signed-rank test on the 6,272
13 paired test samples yields $p = 2.63 \times 10^{-222}$, with
14 4,336 samples showing improved DEM RMSE
15 against 1,936 with higher error, and the gain
16 extends across all reported terrain metrics, sug-
17 gesting that training-data terrain diversity pro-
18 duces a real and broad improvement in transfer
19 to terrain-diverse test regions.

20 These results support the third methodological
21 claim: geographic diversity alone does not guar-
22 antee terrain diversity, and training on terrain-
23 limited data may restrict generalization to more
24 diverse topographic conditions. The compari-
25 son should not be interpreted as a critique of
26 the TfaSR architecture itself, since both models
27 use the same architecture and training config-
28 uration, differing only in their training data.
29 Instead, it shows that the composition of the
30 training dataset is an important factor in DEM
31 SR generalization.

32 6 Discussion

33 DEM SR performance should be interpreted in
34 relation to the evaluation protocol, not only the
35 model architecture. Because nearby DEM tiles
36 often share similar terrain structures, random

37 tile-based testing can give more favourable re-
38 sults than geographically independent testing;
39 random tile splits should therefore not be the
40 only evidence used to claim geographic general-
41 ization.

42 Region-based testing must also be accompanied
43 by terrain-feature analysis: a geographically in-
44 dependent test region can still be artificially easy
45 if dominated by low-relief terrain, or artificially
46 difficult if much more rugged than the training
47 region. Reporting descriptors such as elevation
48 variation, slope, TRI, and TPI clarifies whether
49 reported errors reflect model generalization or
50 test-region difficulty.

51 At the dataset level, geographic coverage should
52 not be treated as equivalent to terrain diver-
53 sity. A dataset collected from multiple areas
54 may still cover a narrow terrain domain, which
55 limits the scope of generalization claims. For
56 this reason, DEM SR evaluation should report
57 performance not only overall but also across
58 terrain-complexity classes. This is especially im-
59 portant for applications where errors in rugged
60 or high-relief terrain can affect downstream ter-
61 rain derivatives and hydrological analyses.

62 This study has several limitations. The exper-
63 iments use TfaSR as the reference model, so
64 numerical results may differ for other architec-
65 tures. The terrain descriptors used here provide
66 a compact representation of terrain complexity
67 but do not capture all possible topographic prop-
68 erties. In addition, the experiments focus on a
69 120 m-to-30 m SRTM-based SR setting; future
70 work should examine other DEM products, scale
71 factors, and geographic regions.

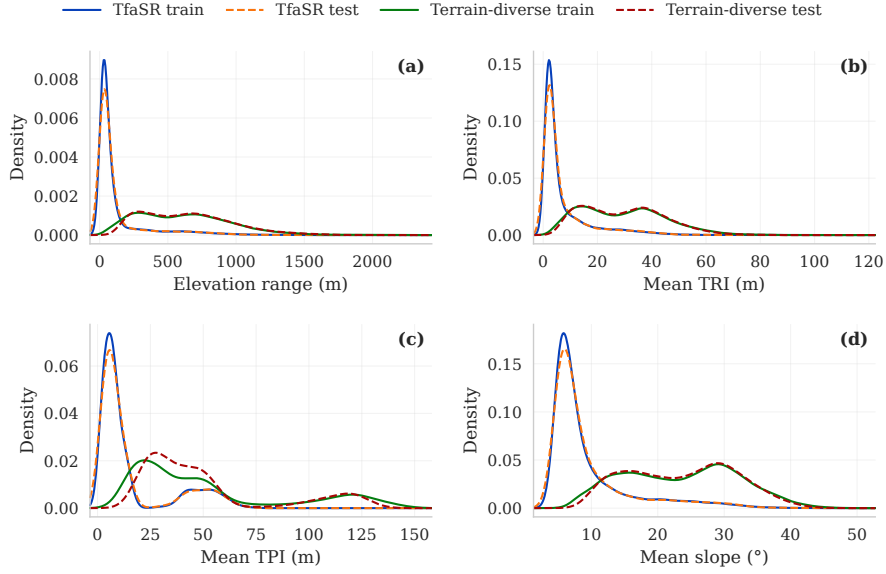


Figure 4: Terrain-feature distributions for Experiment 3. The original TfaSR dataset is concentrated mainly in low-relief terrain, while the terrain-diverse dataset covers a broader range of terrain conditions.

Table 5: Performance of the official pretrained TfaSR model and the TfaSR model trained on the terrain-diverse dataset, both evaluated on the same terrain-diverse test set.

Model	DEM RMSE	TRI RMSE	TPI RMSE	Slope RMSE	Aspect RMSE
Terrain-diverse TfaSR	5.50	7.50	2.96	3.37	57.75
Official pretrained TfaSR	5.69	7.71	3.03	3.48	58.48

7 Conclusion

This paper examined DEM SR evaluation from a methodological perspective. Using controlled experiments with a fixed TfaSR model, we showed that reported performance is affected by random tile-based splitting, test-region terrain representativeness, and dataset-level terrain diversity. These findings suggest that credible DEM SR evaluation should explicitly report spatial partitioning strategies, train-test terrain-feature similarity, dataset-level terrain composition, and class-wise performance across terrain types. Such practices can help ensure that reported improvements reflect meaningful generalization rather than favourable data partitioning or test-region selection.

References

- [1] Ian D. Moore, Richard B. Grayson, and Anthony R. Ladson. Digital terrain modelling: A review of hydrological, geomorphological, and biological applications. *Hydrological Processes*, 5(1):3–30, 1991.
- [2] John P. Wilson and John C. Gallant, editors. *Terrain Analysis: Principles and Applications*. John Wiley & Sons, New York, NY, USA, 2000.
- [3] Paolo Tarolli. High-resolution topography for understanding earth surface processes: Opportunities and challenges. *Geomorphology*, 216:295–312, 2014.
- [4] Tom G. Farr, Paul A. Rosen, Edward Caro, Robert Crippen, Riley Duren, Scott Hensley, Mike Kobrick, Mimi Paller, Ernesto Rodriguez, Ladislav Roth, David Seal, Scott Shaffer, Joanne Shimada, Jeffrey Umland, Marian Werner, Michael Oskin, Douglas Burbank, and Douglas Alsdorf. The shuttle radar topography mission. *Reviews of Geophysics*, 45(2):RG2004, 2007.
- [5] Google Earth Engine Data Catalog. USGS/SRTMGL1_003: NASA SRTM Digital Elevation 30m. https://developers.google.com/earth-engine/datasets/catalog/USGS_SRTMGL1_003, 2024. Accessed: 2026-05-25.
- [6] Bekir Z Demiray, Muhammed Sit, and Ibrahim Demir. D-SRGAN: DEM super-resolution with generative adversarial networks. *SN Computer Science*, 2(1):48, 2021.
- [7] Bekir Z Demiray, Muhammed Sit, and Ibrahim Demir. Dem super-resolution with efficientnetv2. *arXiv preprint arXiv:2109.09661*, 2021.
- [8] Chao Dong, Chen Change Loy, Kaiming He, and

- 1 Xiaoou Tang. Image super-resolution using deep
2 convolutional networks. In *Proceedings of the IEEE*
3 *Conference on Computer Vision and Pattern Recognition*, pages 184–199, 2016.
- 4
- 5 [9] Christian Ledig, Lucas Theis, Ferenc Huszár, Jose
6 Caballero, Andrew Cunningham, Alejandro Acosta,
7 Andrew Aitken, Alykhan Tejani, Johannes Totz, Ze-
8 han Wang, and Wenzhe Shi. Photo-realistic single
9 image super-resolution using a generative adversar-
10 ial network. In *Proceedings of the IEEE Conference*
11 *on Computer Vision and Pattern Recognition*, pages
12 4681–4690, 2017.
- 13 [10] Bee Lim, Sanghyun Son, Heewon Kim, Seungjun
14 Nah, and Kyoung Mu Lee. Enhanced deep resid-
15 ual networks for single image super-resolution. In
16 *Proceedings of the IEEE Conference on Computer*
17 *Vision and Pattern Recognition Workshops*, pages
18 136–144, 2017.
- 19 [11] Jingyun Liang, Jiezhong Cao, Guolei Sun, Kai
20 Zhang, Luc Van Gool, and Radu Timofte. SwinIR:
21 Image restoration using swin transformer. In *Pro-*
22 *ceedings of the IEEE/CVF International Conference*
23 *on Computer Vision Workshops*, pages 1833–1844,
24 2021.
- 25 [12] Yifan Zhang, Wenhao Yu, and Di Zhu. Terrain
26 feature-aware deep learning network for digital el-
27 evation model superresolution. *ISPRS Journal of*
28 *Photogrammetry and Remote Sensing*, 189:143–162,
29 2022.
- 30 [13] Xiaoyi Han, Xiaochuan Ma, Houpu Li, and Zhan-
31 long Chen. A global-information-constrained deep
32 learning network for digital elevation model super-
33 resolution. *Remote Sensing*, 15(2):305, 2023.
- 34 [14] Yi Wang, Shichao Jin, Zekun Yang, Hongcan Guan,
35 Yu Ren, Kai Cheng, Xiaoqian Zhao, Xiaoqiang Liu,
36 Mengxi Chen, Yu Liu, and Qinghua Guo. TTSR:
37 A transformer-based topography neural network
38 for digital elevation model super-resolution. *IEEE*
39 *Transactions on Geoscience and Remote Sensing*,
40 62:1–19, 2024.
- 41 [15] Waldo R. Tobler. A computer movie simulating
42 urban growth in the detroit region. *Economic Ge-*
43 *ography*, 46(sup1):234–240, 1970.
- 44 [16] David R. Roberts, Volker Bahn, Simone Ciuti,
45 Mark S. Boyce, Jane Elith, Gurutzeta Guiller-
46 Arroita, Severin Hauenstein, José J. Lahoz-Monfort,
47 Boris Schröder, Wilfried Thuiller, David I. Warton,
48 Brendan A. Wintle, Florian Hartig, and Carsten F.
49 Dormann. Cross-validation strategies for data
50 with temporal, spatial, hierarchical, or phylogenetic
51 structure. *Ecography*, 40(8):913–929, 2017.
- 52 [17] Pierre Ploton, Frédéric Mortier, Maxime Réjou-
53 Méchain, Nicolas Barbier, Nicolas Picard, Vivien
54 Rossi, Carsten F. Dormann, Gérard Cornu, Guil-
55 laume Viennois, Nicolas Bayol, Alexei Lyapustin,
56 Sylvie Gourlet-Fleury, and Raphaël Pélissier. Spa-
57 tial validation reveals poor predictive performance
58 of large-scale ecological mapping models. *Nature*
59 *Communications*, 11(1):4540, 2020.
- 60 [18] Xiaotong Deng, Weihua Hua, Xiuguo Liu, Siy-
61 ing Chen, Wen Zhang, and Jianchao Duan. D-
62 SRCAGAN : DEM Super-resolution Generative Ad-
63 versarial Network. *IEEE Geoscience and Remote*
64 *Sensing Letters*, pages 1–1, 2022.
- 65 [19] Zhuoxiao Li, Xiaohui Zhu, Shanliang Yao, Yong
66 Yue, Ángel F. García-Fernández, Eng Gee Lim, and
67 Andrew Levers. A large scale digital elevation model
68 super-resolution transformer. *International Journal*
69 *of Applied Earth Observation and Geoinformation*,
70 124:103496, 2023.
- 71 [20] Haijun Feng et al. Information leakage in deep
72 learning-based hyperspectral remote sensing image
73 classification tasks. *Remote Sensing*, 15(15):3793,
74 2023.
- 75 [21] Shawn J. Riley, Stephen D. DeGloria, and Robert
76 Elliot. A terrain ruggedness index that quantifies
77 topographic heterogeneity. *Intermountain Journal*
78 *of Sciences*, 5(1–4):23–27, 1999.
- 79 [22] Andrew D. Weiss. Topographic position and land-
80 forms analysis. In *Poster Presentation, ESRI User*
81 *Conference*, San Diego, CA, USA, 2001.
- 82 [23] Cédric Villani. *Optimal Transport: Old and New*.
83 Springer, 2009.
- 84 [24] George Marsaglia, Wai Wan Tsang, and Jingbo
85 Wang. Evaluating kolmogorov’s distribution. *Jour-*
86 *nal of Statistical Software*, 8(18):1–4, 2003.
- 87 [25] David Lopez-Paz and Maxime Oquab. Revisiting
88 classifier two-sample tests. In *Proceedings of the In-*
89 *ternational Conference on Learning Representations*,
90 2017.