

QualiTeacher: Quality-Conditioned Pseudo-Labeling for Real-World Image Restoration

Fengyang Xiao^{1,*}, Jingjia Feng^{1,*}, Peng Hu², Dingming Zhang¹, Lei Xu³,
Guanyi Qin⁴, Lu Li⁵, Chunming He^{1,†}, and Sina Farsiu¹

¹ Duke University

² Tsinghua University

³ École Polytechnique Fédérale de Lausanne (EPFL)

⁴ National University of Singapore

⁵ Sun Yat-sen University

* Equal Contribution, † Corresponding Author

fengyang.xiao@duke.edu or chunming.he@duke.edu

Abstract. Real-world image restoration (RWIR) is a highly challenging task due to the absence of clean ground-truth images. Many recent methods resort to pseudo-label (PL) supervision, often within a Mean-Teacher (MT) framework, where a teacher network generates targets for a student network. However, these methods face a critical paradox: unconditionally trusting the often imperfect, low-quality PLs forces the student model to learn undesirable artifacts, while discarding them severely limits data diversity and impairs model generalization. In this paper, we propose QualiTeacher, a novel framework that transforms pseudo-label quality from a noisy liability into a conditional supervisory signal. Instead of filtering, QualiTeacher explicitly conditions the student model on the quality of the PLs, estimated by an ensemble of complementary non-reference image quality assessment (NR-IQA) models spanning low-level distortion and semantic-level assessment. This strategy teaches the student network to learn a quality-graded restoration manifold, enabling it to understand what constitutes different quality levels. Consequently, it can not only avoid mimicking artifacts from low-quality labels but also extrapolate to generate results of higher quality than the teacher itself. To ensure the robustness and accuracy of this quality-driven learning, we further enhance the process with a multi-augmentation scheme to diversify the PL quality spectrum, a score-based preference optimization strategy inspired by Direct Preference Optimization (DPO) to enforce a monotonically ordered quality separation, and a cropped consistency loss to prevent adversarial over-optimization (reward hacking) of the IQA models. Experiments on standard RWIR benchmarks demonstrate that QualiTeacher can serve as a plug-and-play strategy to improve the quality of the existing pseudo-labeling framework, establishing a new paradigm for learning from imperfect supervision. Code will be released at <https://github.com/fengyang1399-pixel/QualiTeacher.git>.

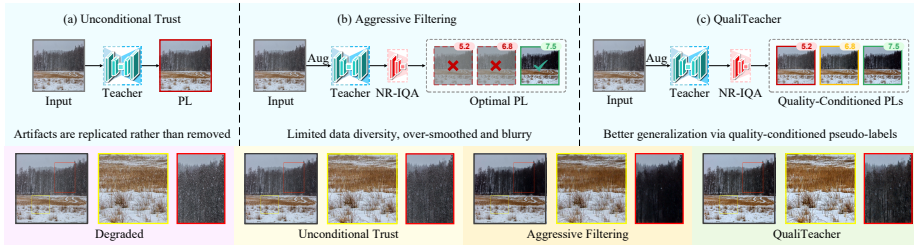


Fig. 1: Comparison of pseudo-label (PL) utilization strategies in the mean-teacher framework. **(a) Unconditional Trust:** PLs are used without filtering, causing degradation artifacts to be replicated by the student. **(b) Aggressive Filtering:** Low-quality PLs are discarded via NR-IQA filtering, yet over-smoothed outputs that receive deceptively high scores survive, introducing blurriness. **(c) QualiTeacher (Ours):** Nearly all PLs are retained (with only extreme outliers discarded), with their quality scores injected as continuous conditioning signals, enabling the student to leverage full data diversity while remaining aware of each sample’s reliability.

1 Introduction

Real-world degradations are complex and spatially varying, encompassing an often unpredictable combination of environmental noise, blur, compression artifacts, and sensor noise [6, 10, 14, 16, 19]. The primary challenge in tackling this real-world image restoration (RWIR) problem lies in the inability to obtain clean, high-quality ground-truth images corresponding to the degraded inputs.

To circumvent this, a popular and effective paradigm is the mean-teacher (MT) framework, which leverages self-supervised or semi-supervised learning. In this setup, a teacher model (often an exponential moving average of the student) generates pseudo-labels (PLs) to supervise the student model. Recent works have focused on refining this process, for instance, by creating “label pools” to store the best historical PLs, with the implicit goal of filtering supervision to include only high-quality targets [10, 17, 21].

Despite their success, they suffer from a fundamental dilemma regarding the quality of the PLs. Owing to the limited capacity of the teacher model, these PLs unavoidably contain residual artifacts or blurred textures. Existing methods thus face a critical paradox, shown in Fig. 1.

This “all-or-nothing” approach severely limits the student’s potential, forcing it to choose between learning from corrupted data or learning from scarce data. In this work, we argue that this dilemma stems from a flawed premise. Instead of treating low-quality PLs solely as noise to be filtered, we should treat them as valuable, graded learning signals to be learned from.

Motivated by this insight, we propose QualiTeacher, a novel framework that transforms this challenge into an opportunity. Our core idea is to make the student model quality-aware. Instead of just learning what to restore, the student learns how the quality of the restoration relates to the underlying content. Specifically, we employ off-the-shelf non-reference Image Quality Assessment (NR-IQA) models to score each PL. This quality score is not used for

filtering; rather, it is embedded (using a continuous embedding, *e.g.*, an MLP or sinusoidal encoding) and injected into the student network as a quality condition.

This quality-conditioned learning strategy achieves two critical goals. First, it allows the student to intelligently discount artifacts from low-quality PLs (*e.g.*, “I am told this target is a 6/10 image, so I will not fully trust its noisy texture”). Second, by training on a spectrum of qualities, the student learns the entire quality-graded restoration manifold. This enables it to extrapolate beyond the teacher’s best performance. Even if the teacher only produces 5-6/10 quality labels, the student learns the vector direction of “quality improvement” and can generate restorations at a quality level exceeding the teacher’s typical output range at inference time by being fed a high-quality score condition.

This ability to actively generate restorations corresponding to specific target scores is the key to breaking the upper bound of the teacher model.

However, merely injecting scores as conditions does not guarantee that the student explicitly learns a monotonic mapping between the condition space and the output quality. To facilitate this, we first propose a multi-augmentation strategy for the teacher. Its goal is not to find the single best label, but to create a rich and diverse curriculum of PLs across a wide quality spectrum, providing the student with varied training data. To address non-uniform quality within an image, we introduce a spatial-aware IQA weighting mechanism to guide the student to focus on high-quality regions. Crucially, to ensure the conditional student produces well-separated outputs proportional to the injected score, we introduce a score-based preference optimization objective inspired by Direct Preference Optimization [40]. By treating NR-IQA scores as automated preference signals, we bypass the need for human annotation and directly optimize the deterministic student model to establish a discriminative, ordered quality manifold. Furthermore, unconstrained optimization toward higher IQA scores may lead to adversarial over-optimization, where the model exploits local statistical shortcuts to inflate scores without genuine perceptual improvement (a phenomenon analogous to reward hacking). To mitigate this, we design a cropped consistency loss that enforces spatial uniformity of quality improvement, thereby ensuring the reliability and robustness of the learned score-quality mapping.

Our main contributions are summarized as follows:

- (1) We propose QualiTeacher for the RWIR task, a novel framework that leverages pseudo-label quality as a condition for deep quality understanding.
- (2) We propose a score-based preference optimization strategy inspired by Direct Preference Optimization to explicitly learn discriminative mappings between score conditions and output quality spaces.
- (3) Experiments verify that QualiTeacher serves as a plug-and-play strategy to improve existing pseudo-labeling frameworks, achieving new SOTAs on several RWIR benchmarks across diverse degradation types.

2 Related Works

2.1 Semi-Supervised Image Restoration

RWIR lacks clean ground-truth for real degraded inputs, motivating semi-supervised paradigms. Existing approaches leverage unpaired data via cycle-consistency [18, 49, 50, 52] or contrastive learning [4, 15, 57], while more recent methods adopt pseudo-label-based strategies [45] where a pre-trained model generates supervision targets. Despite their effectiveness, these methods treat all pseudo-labels uniformly or apply binary filtering, without exploiting the continuous quality spectrum as a learning signal. QualiTeacher fundamentally differs by re-purposing pseudo-label quality as a conditional input rather than a filtering criterion.

2.2 Mean-Teacher Frameworks

The mean-teacher (MT) framework [43] maintains an EMA of the student as the teacher to produce stable pseudo-labels, and has been adapted to desnowing [27], dehazing [10, 11], and other general restoration tasks [38, 54]. Recent works further improve the pipeline through data augmentation [10], consistency regularization [55], and adaptive label selection [33]. However, existing MT-based methods focus on improving pseudo-label reliability through better filtering, while overlooking the quality information embedded in each pseudo-label. QualiTeacher injects the quality score directly into the student, transforming teacher-student learning from unconditional mimicry to quality-conditioned generation.

2.3 NR-IQA in Low-Level Vision

No-reference image quality assessment (NR-IQA) estimates perceptual quality without a reference image, ranging from statistics-based methods like BRISQUE [39] to learning-based approaches such as MUSIQ [26] and CLIP-IQA [44]. In low-level vision, NR-IQA scores have primarily served as post-hoc evaluation metrics [13], with limited use as auxiliary losses [61] or for dataset curation [9, 51]. None of these methods exploits NR-IQA scores as an explicit condition injected into the restoration network. Our work is the first to systematically integrate NR-IQA into every stage of the pipeline, from pseudo-label curation to score-conditioned generation to score-based preference optimization.

2.4 Preference Optimization

Direct Preference Optimization (DPO) [40] bypasses explicit reward modeling by directly optimizing a policy from preference pairs, and has been widely adopted for large language models [53]. In the visual domain, DPO has been applied to diffusion models for text-to-image generation [46], video synthesis [34], and network training [20]. However, preference optimization remains unexplored in the non-generative low-level image restoration tasks. We adapt a score-based preference optimization that bypasses the theoretical obstacle of deterministic models lacking probability outputs, using NR-IQA scores as preference signals and naturally constructing pairs through the score-conditioned strategy.

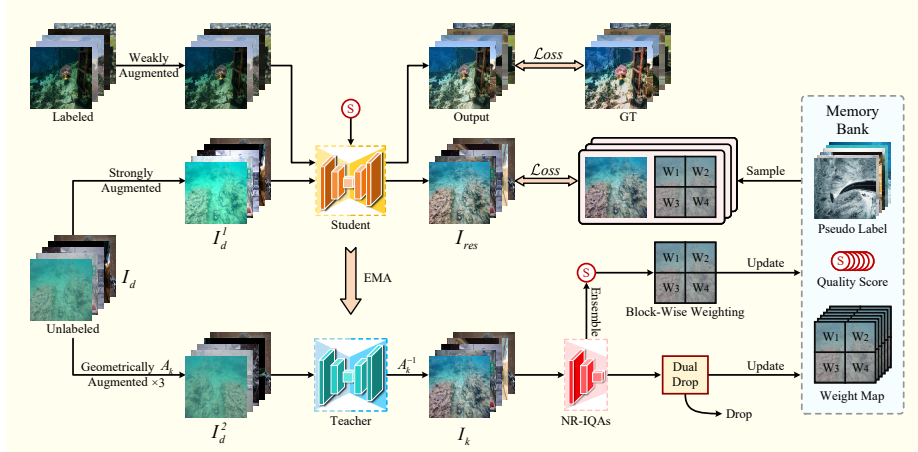


Fig. 2: Overall framework of QualiTeacher. By re-purposing pseudo-label quality as a conditional signal rather than a filtering criterion, the student network can learn a quality-graded manifold and extrapolate beyond the teacher’s capabilities.

3 Methodology

3.1 QualiTeacher

The complete pipeline of QualiTeacher is illustrated in Fig. 2. The framework integrates three synergistic components: (i) a *quality-based pseudo-labeling* module (Sec. 3.2) that curates a diverse spectrum of quality-graded PLs via multi-augmentation, multi-IQA consensus scoring, and dual-drop gating; (ii) a *quality-aware score injection and weighting* module (Sec. 3.3) that embeds IQA scores into the student network as explicit quality conditions with spatial-aware weighting; and (iii) a *quality-driven optimization* module (Sec. 3.4) that incorporates score-based preference optimization and cropped consistency regularization to ensure faithful quality controllability. We elaborate on each component below.

3.2 Quality Based Pseudo-Labeling

Augmentation and Scoring. In the standard MT paradigm, the teacher generates a single pseudo-label per input, yielding a narrow quality distribution that limits supervision diversity. To broaden the quality spectrum of PLs and enhance framework robustness, we introduce a multi-augmentation strategy leveraging geometric transformations to elicit diverse restoration outputs from the teacher. Specifically, given a degraded input I_d , we apply three complementary geometric augmentations (horizontal flip, vertical flip, and 90° rotation), to produce augmented variants $\{\mathcal{A}_k(I_d)\}_{k=1}^3$. Each variant is independently processed by the teacher f_{θ_T} , and the outputs are mapped back to the original coordinate space via corresponding inverse transformations, yielding three PLs:

$$I_k = \mathcal{A}_k^{-1}\left(f_{\theta_T}(\mathcal{A}_k(I_d))\right), \quad k \in \{1, 2, 3\}, \quad (1)$$

where $\mathcal{A}_k^{-1}$ denotes the inverse of the k -th geometric transformation.

Although these augmentations are theoretically lossless, the teacher responds differently to each spatial configuration, naturally producing PLs of varying quality. To evaluate perceptual quality, we employ complementary NR-IQA models across levels: low-level metrics (MUSIQ-KonIQ [26], BRISQUE [39]) capture perceptual distortions, while high-level metrics (CLIP-IQA [44]) leverage vision-language priors for semantic-aware quality estimation. This multi-IQA strategy ensures that quality scores reflect comprehensive and robust assessment rather than the bias of any single metric.

Dual-Drop Mechanism. To ensure multi-IQA consistency and augmentation robustness, we introduce a dual-drop mechanism. In the first drop stage, multiple IQA models jointly evaluate each PL, and the variance across their normalized scores is computed. Given IQA models, $\forall k, m \in \{1, 2, 3\}$, I_k is consistent if:

$$\text{Var}\left(\{\bar{S}_m(I_k)\}_{m=1}^3\right) < \tau_1, \quad \text{Var}\left(\{\bar{S}_m(I_k)\}_{k=1}^3\right) < \tau_2, \quad (2)$$

where $\bar{S}_m$ denotes the min-max normalized score of the m -th IQA model and τ_1 is the consistency threshold. If any PL fails this criterion, the entire augmented set is dropped, enforcing perceptual quality coherence from both low-level and high-level perspectives. Then, each IQA model independently evaluates the full set of PLs, verifying that augmentation yields stable quality estimates, and τ_2 controls the robustness for score variation across geometric transformations.

Gating Mechanism. To ensure high-quality supervision, we maintain a memory bank $\mathcal{B}$ defined to store the PLs. The bank is dynamically updated and retains only the top-3 pseudo-labels from the training history. During training, not all PLs are admitted into the memory bank; only those that pass the dual-drop criterion in each round serve as candidates $\mathcal{C}$ for the gating mechanism. Specifically, the existing PLs in the bank and the candidate PLs are collectively retrieved and re-ranked, from which the new top-3 are selected:

$$\mathcal{B} \leftarrow \text{Top-3}(\mathcal{B} \cup \mathcal{C}, S), \quad (3)$$

where the ranking is determined by the ensemble score S of MUSIQ-KonIQ [26], BRISQUE [39], and CLIP-IQA [44], to :

$$S = 0.4S_{\text{MUSIQ-KonIQ}} + 0.4S'_{\text{BRISQUE}} + 0.2S_{\text{CLIP-IQA}}, \quad (4)$$

where $S'_{\text{BRISQUE}} = (100 - S_{\text{BRISQUE}})$ is the inverted BRISQUE score (converting lower-is-better to higher-is-better), after which all NR-IQA scores are normalized to $[0, 1]$.

3.3 Quality-Aware Score Injection and Weighting

Quality-Aware Score Injection. A central design principle of QualiTeacher is to condition the student network on the perceptual quality of each pseudo-label, transforming restoration from an unconditional mapping into a quality-aware generation process. To achieve this, we inject the aggregated IQA score into the student network at the feature level via a simple yet effective additive mechanism. The conditional student model learns to capture the relationship between PLs and their associated quality scores. During training, when the highest quality score is used as the condition, the PL is set to the highest quality image,

resembling class-conditional generation and enabling the model to learn quality-specific capabilities across diverse levels. Concretely, given the ensemble quality score $S \in \mathbb{R}$ of a pseudo-label, we first map it into a high-dimensional embedding $\mathbf{e}_S \in \mathbb{R}^C$ through a score embedding module:

$$\mathbf{e}_S = \text{MLP}(\gamma(S)), \quad (5)$$

$$\gamma(S) = [\sin(2^0\pi S), \cos(2^0\pi S), \dots, \sin(2^{L-1}\pi S), \cos(2^{L-1}\pi S)]. \quad (6)$$

where C matches the channel dimension of target feature maps and L is the number of frequency bands. The embedding $\mathbf{e}_S$ is then added to intermediate feature representations at selected injection points:

$$\mathbf{F}' = \mathbf{F} + \mathbf{e}_S, \quad (7)$$

where $\mathbf{F}$ denotes the feature at the injection layer, and $\mathbf{e}_S$ is broadcast along spatial dimensions. For instance, in U-Net, injection points are chosen at the encoder-decoder bottleneck, where feature maps carry the most compact and semantically rich representations. More generally, this additive strategy is architecture-friendly with negligible computational overhead, ensuring QualiTeacher serves as a plug-and-play module for various restoration backbones.

Block-wise Evaluation and Weighted Masking. Pseudo-label quality is often spatially non-uniform: certain regions may be well-restored while others retain visible artifacts. To account for such intra-image quality variation, we propose a block-wise evaluation and weighted masking strategy. Specifically, we partition each pseudo-label I_k into a regular grid of 2×2 non-overlapping blocks $\{B_{i,j}\}$, and evaluate the quality of each block independently using the ensemble IQA score Eq. (4). The resulting weight map $\mathbf{W} = \{w_{i,j}\}$ is upsampled to the full resolution and applied element-wise to the pixel-level reconstruction loss, guiding the student to focus on regions where the pseudo-label is most reliable while reducing the influence of artifact-prone areas.

3.4 Quality-Guided Optimization

Score-Based Preference Optimization. In the preceding training stages, the student model has learned to perform restoration conditioned on different quality scores. However, it merely fits the teacher’s pseudo-labels under varying score conditions, without explicitly learning the monotonic mapping between score spaces and output quality. As a result, the learned score spaces may lack sufficient discriminability, leading to outputs generated under different score conditions that may not exhibit quality differences.

To address this, we draw inspiration from Direct Preference Optimization (DPO) [40], originally proposed for aligning language models with human preferences. In our setting, NR-IQA scores serve as an automated preference signal that eliminates the need for human annotation: outputs with higher IQA scores are naturally treated as preferred, while those with lower scores are treated as rejected. Moreover, the score-conditioned student can inherently generate outputs of varying quality by injecting different scores, thereby automatically constructing preference pairs, as shown in Fig. 3

A key theoretical challenge in adapting DPO is that deterministic regression models do not define a probability distribution over outputs, whereas standard DPO requires computing log-probabilities. We bypass this by operating in the output quality space. Specifically, our score-conditioned architecture naturally produces distinct outputs y^h and y^l by injecting different quality scores, constructing preference pairs without stochastic sampling or likelihood estimation. The NR-IQA evaluator then acts as a deterministic reward function mapping each output to a scalar quality measure, replacing log-probabilities in the original DPO formulation. This shifts the optimization target from “making the preferred output more likely” to “making the high-score output measurably better,” enabling preference optimization in a fully deterministic pipeline.

Unlike standard DPO, which optimizes a single preference direction, our formulation establishes a score-conditioned preference ordering: the objective is not merely to prefer one output over another, but to ensure that restoration quality varies monotonically with the injected score condition, forming a well-separated quality manifold. Eq. (8) explicitly forces the student to learn score-space-specific features under different score conditions, producing outputs with clearly distinguishable and ordered quality levels.

$$\begin{aligned} \mathcal{L}_{pref} = & \underbrace{-\log \sigma\left(\beta(S_{stu}(y^h, x) - S_{stu}(y^l, x) - \delta)\right)}_{\mathcal{L}_1} \\ & - \underbrace{\log \sigma\left(\beta(S_{stu}(y^h, x) - S_{tea}(I_k^h, x))\right)}_{\mathcal{L}_2} + \underbrace{\frac{1}{2} \sum_{n=1}^N \|W_n - W_n^{\text{init}}\|_F^2}_{\mathcal{L}_{reg}}, \end{aligned} \quad (8)$$

where $S_{stu}(y, x)$ denotes the IQA-evaluated quality of the student’s output y given input x , $S_{tea}(I_k^h, x)$ denotes the quality of the teacher’s best pseudo-label, y^h and y^l are the student’s outputs conditioned on high and low quality scores respectively, $\sigma(\cdot)$ is the sigmoid function, β is a temperature controlling preference sharpness, and δ is a margin specifying the minimum desired quality gap between score spaces. W_n is the current weights of the n th score injection convolution layer. $\|\cdot\|_F^2$ denotes the squared Frobenius norm.

Loss 1: $\mathcal{L}_1$. This term enforces inter-space separation: it encourages the student to produce outputs with clearly distinguishable quality when conditioned on different scores, by penalizing cases where the quality gap between y^h and y^l falls below the margin δ . This operates within the student’s output space, ensuring that the model learns distinct score spaces and corresponding features.

Loss 2: $\mathcal{L}_2$. While $\mathcal{L}_1$ enlarges the gap between score spaces, it does not specify the *direction* of separation, that is, the model may achieve a large gap by degrading the low-score output rather than improving the high-score one. To prevent this, $\mathcal{L}_2$ introduces a soft lower bound on the student’s high-score output, anchoring the separation upward: the student must improve its best output rather than merely worsening its worst. Together, $\mathcal{L}_1$ and $\mathcal{L}_2$ are complementary: $\mathcal{L}_1$ ensures different score conditions produce distinguishable outputs, while $\mathcal{L}_2$ guarantees this separation progressively pushes overall quality upward.

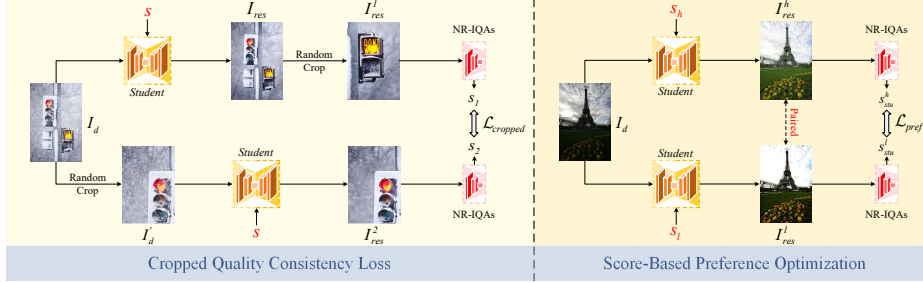


Fig. 3: Overview of the quality-driven optimization strategy in QualiTeacher. Left: Cropped Quality Consistency Loss. Right: Score-Based Preference Optimization.

Loss 3: $\mathcal{L}_{reg}$. Directly optimizing the preference objective can interfere with underlying restoration capabilities; that is, the image may score higher with the IQA model yet visually contain artifacts. We design $\mathcal{L}_{reg}$ to anchor the injection layer weights near their initial values, ensuring that preference optimization improves quality only through slight, quality-relevant modifications without significantly altering the overall network behavior.

Cropped Quality Consistency Loss. Using NR-IQA scores as optimization objectives is promising but inherently unstable: without a reference, the model may generate outputs that deceive IQA models into assigning high scores while exhibiting low quality, a phenomenon analogous to adversarial examples and reward hacking in large-scale reward learning systems. To mitigate this, as shown in Fig. 3, we propose a cropped quality consistency loss that enforces spatial robustness of the quality evaluation. The key idea is to compare two processing orders: (i) restore the full image, then randomly crop a quarter of the restored image and evaluate to obtain S_1 via Eq. (4); (ii) randomly crop the degraded input first, then restore and evaluate to obtain S_2 . The loss is defined as:

$$\mathcal{L}_{cropped} = \|S_1 - S_2\|_1. \quad (9)$$

If the model exploits local statistical shortcuts, such as inserting imperceptible artifacts, to inflate global IQA scores, these artifacts are typically spatially unstable and will manifest as discrepancies between S_1 and S_2 . Furthermore, NR-IQA models often exhibit scale-dependent behavior, producing reliable assessments on full images but inconsistent evaluations on local patches. By enforcing consistency between the two processing orders, this loss prevents the model from exploiting scale-dependent biases of NR-IQA models, ensuring that quality improvements are genuine and spatially uniform.

3.5 Training Strategy

The teacher network parameters θ_T are updated as an exponential moving average (EMA) of the student parameters θ_S (with $\alpha = 0.998$):

$$\theta_T \leftarrow \alpha \theta_T + (1 - \alpha) \theta_S. \quad (10)$$

In the training stage, we train the student model with score injection, score-based preference optimization, and cropped consistency loss, enabling it to acquire quality awareness, restoration ability, and structural fidelity, learning how

to map degraded inputs to a different manifold. The objective is:

$$\mathcal{L} = \mathcal{L}_{rec} + \mathcal{L}_{per} + \mathcal{L}_{pref} + \mathcal{L}_{cropped}, \quad (11)$$

where $\mathcal{L}_{rec} = \mathbf{W} \odot \|f_{\theta_S}(I_d, S) - I_k^*\|_1$ is the pixel-level ℓ_1 reconstruction loss between the student output and the corresponding pseudo-label I_k^* selected from the memory bank, and $\mathcal{L}_{per}$ is a perceptual loss based on VGG features. This formulation ensures that the resulting model is both quality-discriminative and restoration-faithful, balancing perceptual preference with pixel-level accuracy.

4 Experiment

4.1 Implementation Details

QualiTeacher is implemented in PyTorch and trained on two NVIDIA A6000 GPUs using AdamW with momentum terms of (0.9, 0.999). The learning rate is set to 5×10^{-5} with only 10K iterations. Since pseudo-labels with scores above 7 are relatively scarce, training a separate score bin for each integer level would lead to insufficient samples and unstable optimization. To address this, we merge all pseudo-labels with quality scores exceeding 7 into a unified high-quality space. Consequently, the maximum guidance score during inference is set to 7.

During the experiment, we use MUSIQ [26], BRISQUE [39], and CLIP-IQA [44] for training. Since no ground-truth reference images are available for real-world degraded inputs, we employ a comprehensive set of NR-IQA metrics for evaluation across all tasks. These metrics span CNN-based (CNNIQA [25], HYPERIQA [41], MANIQA [58]), Transformer-based (TOPIQ [5], ARNIQA [2]), and vision-language model-based (Q-Align [48], Compare2Score [63], LIQE [62]) methods, where higher values indicate better perceptual quality. In addition, we include task-specific metrics where applicable: FADE [7] for dehazing, which measures residual fog density (lower is better); MFRQA [31] for LLIE, which evaluates brightness, color, structure, and naturalness (higher is better); and URANKER [12] for UIE (higher is better). The specific subset of metrics used for each task is detailed in Tab. 1.

4.2 Comparative Evaluation

Desnowing. With NAFNet as the backbone, we train on paired data from Snow100K [36] and unpaired data from RealSnow 10K-Train [64], evaluating on RealSnow 10K-Test [64] against SnowMaster [27] and Colabator. To further probe generalizability, we substitute the backbone with a SOTA model HistoFormer [42] and benchmark against Colabator [10], alongside the end-to-end semi-supervised method SemiDDM-Weather [38]. As reported in Tab. 1 and Fig. 4, QualiTeacher consistently outperforms Colabator and existing semi-supervised methods across all metrics. Furthermore, it delivers visually cleaner outputs with sharper textures and fewer residual snow artifacts, confirming the generalization of our quality-prior conditioning strategy to the desnowing task.

Deraining. For deraining, we build upon the SOTA weather restoration model HistoFormer [42] and compare QualiTeacher with Colabator [10], as well as

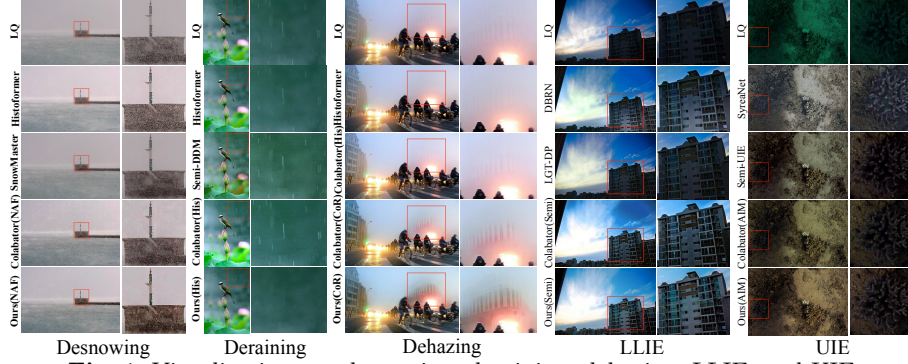


Fig. 4: Visualizations on desnowing, deraining, dehazing, LLIE, and UIE.

Table 1: NR-IQA comparisons on Desnowing, Deraining, Dehazing, LLIE, and UIE.

Desnowing								
Methods	CNNIQA↑	HyperIQA↑	ManIQA↑	Topiq↑	Arniqa↑	Q-Align↑	C2score↑	LIQE↑
SemiDDM-Weather [38]	0.636	0.423	0.342	0.442	0.567	3.586	55.068	2.497
SnowMaster (NAFNet) [27]	0.670	0.501	0.349	0.499	0.612	3.836	55.327	2.768
NAFNet + Colabator [10]	0.628	0.470	0.342	0.473	0.620	3.936	55.389	2.667
NAFNet + QualiTeacher	0.675	0.559	0.362	0.540	0.628	4.030	55.446	2.883
HistoFormer [42]	0.615	0.444	0.346	0.455	0.585	3.827	55.266	2.610
HistoFormer + Colabator [10]	0.642	0.480	0.365	0.477	0.586	3.712	55.075	2.297
HistoFormer + QualiTeacher	0.674	0.530	0.368	0.496	0.592	4.018	55.314	2.429
Deraining								
Methods	CNNIQA↑	HyperIQA↑	ManIQA↑	Topiq↑	Arniqa↑	Q-Align↑	C2score↑	LIQE↑
SSID-KD [8]	0.580	0.423	0.300	0.427	0.579	3.641	54.790	2.475
MOSS [23]	0.559	0.415	0.287	0.416	0.597	3.720	54.803	2.467
SemiDDM-Weather [38]	0.620	0.463	0.294	0.438	0.583	3.583	54.712	2.688
HistoFormer [42]	0.584	0.418	0.299	0.427	0.576	3.734	54.810	2.556
HistoFormer + Colabator [10]	0.690	0.545	0.308	0.520	0.609	3.690	54.864	2.680
HistoFormer + QualiTeacher	0.694	0.547	0.308	0.508	0.613	3.783	54.911	2.539
Dehazing								
Methods	CNNIQA↑	HyperIQA↑	Topiq↑	Arniqa↑	Q-Align↑	C2score↑	LIQE↑	FADE ↓
Colabator (CORUN) [10]	0.686	0.602	0.599	0.718	3.570	55.272	2.946	0.824
CORUN + QualiTeacher	0.706	0.621	0.601	0.724	3.628	55.322	3.317	0.590
HistoFormer [42]	0.543	0.415	0.402	0.582	3.050	54.291	1.989	2.499
HistoFormer + Colabator [10]	0.690	0.579	0.547	0.658	3.031	54.363	2.473	1.216
HistoFormer + QualiTeacher	0.702	0.550	0.553	0.665	3.270	54.738	2.460	1.145
Low-Light Image Enhancement								
Methods	CNNIQA↑	ManIQA↑	Topiq↑	Arniqa↑	Q-Align↑	C2score↑	LIQE↑	MFRQA↑
DRBN [59]	0.564	0.336	0.483	0.652	3.780	55.129	3.136	56.198
LMT-GP (SemiLL) [60]	0.561	0.387	0.720	0.673	3.655	55.387	3.095	58.866
SemiLL + Colabator [10]	0.614	0.372	0.724	0.671	3.730	55.103	3.320	58.138
SemiLL + QualiTeacher	0.626	0.411	0.733	0.692	3.804	55.501	3.376	59.372
CIDNet [56]	0.554	0.354	0.480	0.618	3.715	54.944	3.049	58.860
CIDNet + Colabator [10]	0.532	0.367	0.439	0.609	3.571	54.967	2.929	59.319
CIDNet + QualiTeacher	0.669	0.428	0.542	0.657	3.741	55.400	3.322	59.341
Underwater Image Enhancement								
Methods	CNNIQA↑	ManIQA↑	Topiq↑	Arniqa↑	Q-Align↑	C2score↑	LIQE↑	URANKER↑
SyreaNet [47]	0.449	0.195	0.320	0.513	3.630	55.034	1.658	1.591
Semi-UIR (AIM-Net) [24]	0.438	0.202	0.323	0.474	3.706	55.322	1.642	2.215
AIM-Net + Colabator [10]	0.491	0.244	0.365	0.542	3.825	55.510	1.729	2.374
AIM-Net + QualiTeacher	0.515	0.263	0.376	0.548	3.835	55.639	1.804	2.467

existing semi-supervised baselines SSID-KD [8], MOSS [23], and SemiDDM-Weather [38]. Both the quantitative results in Tab. 1 and the visual comparisons in Fig. 4 evaluated on Real 3000 [35] confirm a clear margin: QualiTeacher

Table 2: Break down ablation.

ID	Score Condition	Preference Optimization	Cropped Consistency	Drop Strategy	ManIQA $\uparrow$	Arniqa $\uparrow$	HyperIQA $\uparrow$	Q-Align $\uparrow$	C2score $\uparrow$	LIQE $\uparrow$
(a)	×	×	×	×	0.329	0.574	0.452	3.850	55.309	2.559
(b)	✓	×	×	×	0.341	0.601	0.510	4.001	55.399	2.719
(c)	✓	✓	×	×	0.352	0.613	0.535	4.021	55.431	2.805
(d)	✓	✓	✓	×	0.358	0.620	0.543	4.027	55.442	2.860
(e)	✓	✓	✓	✓	0.362	0.628	0.559	4.030	55.446	2.883

Table 3: User study on Desnowing, Dehazing, and LLIE.

(a) Desnowing			(b) Dehazing			(c) LLIE		
Methods	Score	Rank	Methods	Score	Rank	Methods	Score	Rank
SemiDDM-Weather [38]	2.92	4	Colabator (CORUN) [10]	3.35	2	DRBN [59]	3.67	4
SnowMaster (NAFNet) [27]	3.42	2	CORUN + QualiTeacher	3.62	1	LMT-GP (SemiLL) [60]	3.85	2
NAFNet + Colabator [10]	3.23	3	HistoFormer + Colabator [10]	2.27	4	SemiLL + Colabator [10]	3.71	3
NAFNet + QualiTeacher	3.58	1	HistoFormer + QualiTeacher	2.81	3	SemiLL + QualiTeacher	3.96	1

significantly outperforms all competing methods on every metric, while better preserving fine-grained textures without introducing over-smoothing artifacts.

Dehazing. Following CORUN [10], dehazing performance is assessed on RTTS [30]. We additionally swap the backbone to HistoFormer [42] to test generalizability. Tab. 1 reveals that QualiTeacher achieves the strongest results across all metrics, with particularly pronounced gains on FADE [7], underscoring the capacity of our semi-supervised framework to suppress residual haze.

Low-Light Image Enhancement. We adopt SemiLL as the default backbone following LMT-GP [60], retaining its original training configuration with VE-LOL [32] for training and DICM [28] for inference. Comparisons are conducted against Colabator [10] and DRBN [59], which do not support backbone replacement. We further substitute the backbone with CIDNet [56] to validate generalizability. As reported in Tab. 1, QualiTeacher leads across all eight metrics. The visualizations in Fig. 4 reinforce this advantage: our framework produces more stable and robust enhancements, with notably improved brightness, contrast, and structural fidelity, highlighting the practical value of quality-aware guidance in low-light scenarios.

Underwater Image Enhancement. One backbone architecture is employed: AIM-Net from Semi-UIR [24], following the original training configurations of Semi-UIR [24]. Inference is performed on Seathru [3], with Colabator [10], Semi-UIR [24] and SyreaNet [47]. Tab. 1 and Fig. 4 show that QualiTeacher establishes new state-of-the-art results across all metrics, producing underwater images with markedly enhanced color fidelity and finer structural details, demonstrating the effectiveness of quality-aware guidance in this challenging domain.

4.3 Ablation Study

We use NAFNet as the backbone, the RealSnow 10K-Test [64] as the inference dataset, and conduct ablation studies on the desnowing task. For space limitations, only part ablations are included; please see supplementary material for more results, such as the IQA ensemble strategy and score injection mechanism.

Quality Prior Conditional Approach. To steer the student model toward optimal restoration quality, we adopt a quality score of 7 as the conditioning signal. The progression from (a) to (b) in Tab. 2 confirms that integrating this quality prior into the student network yields consistent improvements. As shown

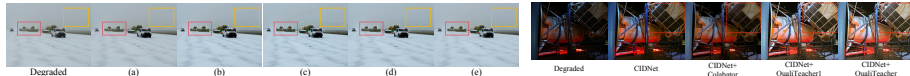


Fig. 5: Visualization of breakdown ablation, where (a), (b), (c), (d), and (e) are consistent with those in Tab. 2.



Fig. 6: Visualization on LLIE with changed IQA selection (MANIQA, TOPIQ, QualiClip).

in Fig. 5, outputs from (b) present noticeably finer details and snowy artifact removal over (a), validating the benefit of incorporating perceptual priors.

Effectiveness of Preference Optimization. Building on the quality prior, preference optimization further elevates the student model’s output, as reflected in (b) $\rightarrow$ (c) of Tab. 2. As Fig. 5 shows, quality score conditioning steers the student toward producing higher-quality outputs. However, this IQA-driven optimization can sometimes be a double-edged sword, as it may inadvertently introduce grid-like background artifacts while pursuing higher quality scores.

Effectiveness of Cropped Consistency Loss. This component targets the over-optimization problem commonly encountered in continuous score spaces. As shown by (c) $\rightarrow$ (d) in Tab. 2, the cropped consistency loss strengthens reconstruction quality while preserving stability and effectively addressing overoptimization. Visual comparisons in Fig. 5 echo this observation, revealing refined fine-grained details, enhanced overall perceptual realism, and removed grid-like artifacts in the background.

Effectiveness of Dual-Drop Mechanism. The dual-drop mechanism yields a notable performance improvement, as evidenced by the comparison between (d) and (e) in Tab. 2. Visual results in Fig. 5 further corroborate this finding, where outputs from (e) exhibit sharper edges, richer textures, and aesthetic improvement compared to those of (d). This demonstrates that the dual-drop strategy enhances both the robustness and the output quality of the framework.

4.4 Further Analysis and Extended Applications

User Study. We conduct a user study on desnowing (RealSnow 10K [64]), dehazing (RTTS [30]), and LLIE (DICM [28]). Fifteen participants rated restored images on a 0-5 scale based on visibility, artifact removal, and naturalness. Each degraded image and its restored counterpart were displayed side by side, with method identities hidden. As shown in Tab. 3, our method consistently receives the highest ratings, confirming its superiority in perceptual quality.

IQA Selection. Our IQA selection follows a complementary-level principle, combining low-level and high-level metrics for comprehensive quality supervision. To verify that performance gains stem from our quality-prior mechanism rather than specific IQA choices, we substitute three alternative NR-IQA models (MANIQA [58], TOPIQ [5], and QualiClip [1]) under the same principle and retrain the framework to get QualiTeacher1. As shown in Tab. 4 and Fig. 6, evaluation on LLIE with eight metrics under identical settings still achieves state-of-the-art performance, confirming that improvement is attributed to our strategy design rather than particular IQA engineering selection as well as verifying the generalizability of our quality-prior mechanism.

Table 4: Ablation on the IQA selection. QualiTeacher with alternative NR-IQA metrics (MANIQA, TOPIQ, QualiClip) on LLIE.

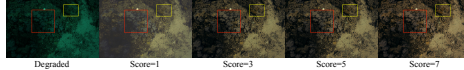
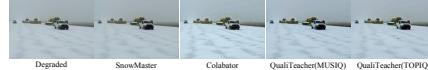
Methods	CNNIQA↑	ManIQA↑	Topiq↑	Arniqa↑	Q-Align↑	C2score↑	LIQE↑	MFRQA↑
CIDNet [56]	0.554	0.354	0.480	0.618	3.715	54.944	3.049	58.860
CIDNet + Colabator [10]	0.532	0.367	0.439	0.609	3.571	54.967	2.929	59.319
CIDNet + QualiTeacher1	0.660	0.432	0.530	0.661	3.643	55.454	3.411	59.675
CIDNet + QualiTeacher(Ours)	0.669	0.428	0.542	0.657	3.741	55.400	3.322	59.341

Table 5: Low-light image detection on *ExDark*.

Methods (AP)	Bicycle	Boat	Bottle	Bus	Car	Cat	Chair	Cup	Dog	Motor	People	Table	Mean
DRBN [59]	81.8	81.1	73.5	90.2	79.6	72.6	65.1	69.9	78.2	71.2	76.7	62.7	75.2
LMT-GP(SemiLL) [60]	80.9	82.7	76.2	91.2	80.2	75.1	72.1	70.6	77.8	71.1	78.8	65.8	76.9
SemiLL + Colabator [10]	79.6	82.0	75.7	89.2	80.8	69.5	70.0	70.8	79.9	71.6	76.8	65.6	76.0
CIDNet + Colabator [10]	80.7	79.8	76.0	86.6	81.2	75.2	71.9	71.9	79.6	72.1	78.3	67.3	76.7
SemiLL + QualiTeacher	82.6	83.3	78.2	91.2	82.1	70.9	74.6	72.0	80.7	74.2	80.5	67.9	78.2
CIDNet + QualiTeacher	82.9	83.0	76.4	92.5	82.5	77.8	73.3	73.5	82.8	74.0	81.4	68.7	79.1

Table 6: Robust analysis for NR-IQA metrics. We replace MUSIQ with TOPIQ.

Methods (Task: Desnoising)	CNNIQA↑	HyperIQA↑	ManIQA↑	Topiq↑	Arniqa↑	Q-Align↑	C2score↑	LIQE↑
SnowMaster(NAFNet) [27]	0.670	0.501	0.349	0.499	0.612	3.836	55.327	2.768
NAFNet + Colabator [10]	0.628	0.470	0.342	0.473	0.620	3.936	55.389	2.667
NAFNet + QualiTeacher(Topiq)	0.671	0.538	0.361	0.540	0.614	3.999	55.398	2.826
NAFNet + QualiTeacher	0.675	0.559	0.362	0.540	0.628	4.030	55.446	2.883

**Fig. 7:** Score controllability.**Fig. 8:** Failure cases analysis.

Score Controllability Analysis. We show that the condition score directly influences restoration quality at test time. As illustrated in Fig. 7, by varying the injected score at inference, the model generates outputs of corresponding quality, enabling continuous control over perceptual quality. In practice, we set the condition to the maximum score to achieve optimal restoration results. Notably, quality controllability proves particularly beneficial for UIE, where higher score guidance yields more visually balanced and realistic color tones.

Benefits for Downstream Tasks. To validate that improving image quality facilitates downstream tasks, we evaluate enhancement methods on low-light object detection. We enhance low-light images from ExDark [37] using various algorithms and feed the results into YOLOv13 [29] for detection. As reported in Tab. 5, QualiTeacher consistently boosts detection performance.

5 Limitations

As shown in Fig. 8, QualiTeacher occasionally produces grid-like artifacts in a very small number of images, while Colabator exhibits them in a considerably larger portion of its results. Since reducing training iterations only mitigates but does not eliminate these artifacts, we hypothesize that the root cause lies in the MUSIQ-KonIQ metric employed by both methods. Although our framework adopts an ensemble strategy to minimize single-metric bias, MUSIQ is trained on KonIQ [22], a photography-oriented dataset whose high-quality references tend to contain sharp, repetitive texture patterns. Consequently, MUSIQ may assign higher scores to outputs with grid-like structures, inadvertently guiding the

model toward such artifacts. To verify this, we replace Musiq with Topiq [5]. As shown in Tab. 6 and Fig. 8, the grid artifacts vanish, confirming our conjecture. Notably, QualiTeacher (TOPIQ) still achieves SOTA performance, demonstrating our framework’s robustness to the choice of IQA metrics.

6 Conclusion

In this paper, we propose QualiTeacher, a novel framework that reframes pseudo-label quality as a conditional signal instead of a filter. This quality-conditioned approach enables our student model to learn a complete restoration manifold, which critically allows it to extrapolate beyond the teacher’s quality ceiling and achieve state-of-the-art results. Our work, proven robust across various NR-IQA models, establishes a new paradigm for learning from imperfect supervision and opens new avenues for flexible restoration strategies.

7 Acknowledgement

This work was supported in part by the Foundation Fighting Blindness (BR-CL-0621-0812-DUKE and PPA-1224-0890-DUKE) and a Research to Prevent Blindness Unrestricted Grant to Duke University.

References

1. Agnolucci, L., Galteri, L., Bertini, M.: Quality-aware image-text alignment for opinion-unaware image quality assessment. arXiv preprint arXiv:2403.11176 (2024) 13
2. Agnolucci, L., Galteri, L., Bertini, M., Del Bimbo, A.: Arniqa: Learning distortion manifold for image quality assessment. In: WACV. pp. 189–198 (2024) 10
3. Akkaynak, D., Treibitz, T.: Sea-thru: A method for removing water from underwater images. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 1682–1691 (2019) 12
4. Basak, H., Yin, Z.: Pseudo-label guided contrastive learning for semi-supervised medical image segmentation. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 19786–19797 (2023) 4
5. Chen, C., Mo, J., Hou, J., Wu, H., Liao, L., Sun, W., Yan, Q., Lin, W.: Topiq: A top-down approach from semantics to distortions for image quality assessment. IEEE Transactions on Image Processing 33, 2404–2418 (2024) 10, 13, 15
6. Chen, Z., Zhang, Y., Liu, D., Xia, B., Gu, J., Kong, L., Yuan, X.: Hierarchical integration diffusion model for realistic image deblurring. In: NeurIPS (2023) 2
7. Choi, L.K., You, J., Bovik, A.C.: Referenceless prediction of perceptual fog density and perceptual image defogging. IEEE Transactions on Image Processing 24(11), 3888–3901 (2015) 10, 12
8. Cui, X., Wang, C., Ren, D., Chen, Y., Zhu, P.: Semi-supervised image deraining using knowledge distillation. IEEE Transactions on Circuits and Systems for Video Technology 32(12), 8327–8341 (2022) 11

9. Dai, X., Hou, J., Ma, C.Y., Tsai, S., Wang, J., Wang, R., Zhang, P., Vandenhende, S., Wang, X., Dubey, A., et al.: Emu: Enhancing image generation models using photogenic needles in a haystack. arXiv preprint arXiv:2309.15807 (2023) [4](#)
10. Fang, C., He, C., Xiao, F., Zhang, Y., Tang, L., Zhang, Y., Li, K., Li, X.: Real-world image dehazing with coherence-based label generator and cooperative unfolding network. NeurIPS (2024) [2](#), [4](#), [10](#), [11](#), [12](#), [14](#)
11. Fang, C., He, C., Zhang, Y., Chen, C., Zhu, C., Tang, L., Li, X.: Prism: Rethinking scattered atmosphere reconstruction as a unified understanding and generation model for real-world dehazing. arXiv preprint arXiv:2604.07048 (2026) [4](#)
12. Guo, C., Wu, R., Jin, X., Han, L., Zhang, W., Chai, Z., Li, C.: Underwater ranker: Learn which is better and how to be better. In: AAAI. vol. 37, pp. 702–709 (2023) [10](#)
13. He, C., Fang, C., Zhang, Y., Li, K., Tang, L., You, C., Xiao, F., Guo, Z., Li, X.: Reti-diff: Illumination degradation image restoration with retinex-based latent diffusion model. arXiv preprint arXiv:2311.11638 (2023) [4](#)
14. He, C., Li, K., Xu, G., Zhang, Y., Hu, R., Guo, Z., Li, X.: Degradation-resistant unfolding network for heterogeneous image fusion. In: ICCV. pp. 12611–12621 (2023) [2](#)
15. He, C., Shen, Y., Fang, C., Xiao, F., Tang, L., Zhang, Y., Zuo, W., Guo, Z., Li, X.: Diffusion models in low-level vision: A survey. arXiv preprint arXiv:2406.11138 (2024) [4](#)
16. He, C., Xiao, F., Zhang, R., Fang, C., Fan, D.P., Farsiu, S.: Reversible unfolding network for concealed visual perception with generative refinement. arXiv preprint arXiv:2508.15027 (2025) [2](#)
17. He, C., Zhang, R., Chen, Z., Yang, B., Fang, C., Lin, Y., Xiao, F., Farsiu, S.: Unfoldldm: Deep unfolding-based blind image restoration with latent diffusion priors. arXiv preprint arXiv:2511.18152 (2025) [2](#)
18. He, C., Zhang, R., Xiao, F., Fang, C., Tang, L., Zhang, Y., Farsiu, S.: Unfoldir: Rethinking deep unfolding network in illumination degradation image restoration. CVPR (2026) [4](#)
19. He, C., Zhang, R., Xiao, F., Fang, C., Tang, L., Zhang, Y., Kong, L., Fan, D.P., Li, K., Farsiu, S.: Run: Reversible unfolding network for concealed object segmentation. arXiv preprint arXiv:2501.18783 (2025) [2](#)
20. He, C., Zhang, R., Xiao, F., Zhang, D., Cao, Z., Farsiu, S.: Refining context-entangled content segmentation via curriculum selection and anti-curriculum promotion. In: ICML (2026) [4](#)
21. He, C., Zhang, R., Zhang, D., Xiao, F., Fan, D.P., Farsiu, S.: Nested unfolding network for real-world concealed object segmentation. ECCV (2026) [2](#)
22. Hosu, V., Lin, H., Sziranyi, T., Saupe, D.: KonIQ-10k: An ecologically valid database for deep learning of blind image quality assessment. IEEE Transactions on Image Processing **29**, 4041–4056 (2020) [14](#)
23. Huang, H., Yu, A., He, R.: Memory oriented transfer learning for semi-supervised image deraining. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 7732–7741 (2021) [11](#)
24. Huang, S., Wang, K., Liu, H., Chen, J., Li, Y.: Contrastive semi-supervised learning for underwater image restoration via reliable bank. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 18145–18155 (2023) [11](#), [12](#)
25. Kang, L., Ye, P., Li, Y., Doermann, D.: Convolutional neural networks for no-reference image quality assessment. In: CVPR. pp. 1733–1740 (2014) [10](#)

26. Ke, J., Wang, Q., Wang, Y., Milanfar, P., Yang, F.: Musiq: Multi-scale image quality transformer. In: Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV). pp. 5148–5157 (October 2021) [4](#), [6](#), [10](#)
27. Lai, J., Chen, S., Lin, Y., Ye, T., Liu, Y., Fei, S., Xing, Z., Wu, H., Wang, W., Zhu, L.: Snowmaster: Comprehensive real-world image desnowing via mllm with multi-model feedback optimization. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 4302–4312 (2025) [4](#), [10](#), [11](#), [12](#), [14](#)
28. Lee, C., Lee, C., Kim, C.S.: Contrast enhancement based on layered difference representation of 2d histograms. *IEEE Trans. Image Process.* **22**(12), 5372–5384 (2013) [12](#), [13](#)
29. Lei, M., Li, S., Wu, Y., Hu, H., Zhou, Y., Zheng, X., Ding, G., Du, S., Wu, Z., Gao, Y.: Yolov13: Real-time object detection with hypergraph-enhanced adaptive visual perception. *arXiv preprint arXiv:2506.17733* (2025) [14](#)
30. Li, B., Ren, W., Fu, D., Tao, D., Feng, D., Zeng, W., Wang, Z.: Benchmarking single-image dehazing and beyond. *IEEE transactions on image processing* **28**(1), 492–505 (2018) [12](#), [13](#)
31. Lin, W., Wu, Y., Xu, L., Chen, W., Zhao, T., Wei, H.: No-reference quality assessment for low-light image enhancement: Subjective and objective methods. *Displays* **78**, 102430 (2023) [10](#)
32. Liu, J., Xu, D., Yang, W., Fan, M., Huang, H.: Benchmarking low-light image enhancement and beyond. *International Journal of Computer Vision* **129**(4), 1153–1184 (2021) [12](#)
33. Liu, L., Zhang, B., Zhang, J., Zhang, W., Gan, Z., Tian, G., Zhu, W., Wang, Y., Wang, C.: Mixteacher: Mining promising labels with mixed scale teacher for semi-supervised object detection. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 7370–7379 (2023) [4](#)
34. Liu, R., Wu, H., Zheng, Z., Wei, C., He, Y., Pi, R., Chen, Q.: Videodpo: Omni-preference alignment for video diffusion generation. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 8009–8019 (2025) [4](#)
35. Liu, Y., Yue, Z., Pan, J., Su, Z.: Unpaired learning for deep image deraining with rain direction regularizer. In: Proceedings of the IEEE/CVF international conference on computer vision. pp. 4753–4761 (2021) [11](#)
36. Liu, Y.F., Jaw, D.W., Huang, S.C., Hwang, J.N.: Desnownet: Context-aware deep network for snow removal. *IEEE Transactions on Image Processing* **27**(6), 3064–3073 (2018) [10](#)
37. Loh, Y.P., Chan, C.S.: Getting to know low-light images with the exclusively dark dataset. *Computer Vision and Image Understanding* **178**, 30–42 (2019) [14](#)
38. Long, F., Su, W., Li, Z., Cai, L., Li, M., Wang, Y.G., Cao, X.: Semiddm-weather: A semi-supervised learning framework for all-in-one adverse weather removal. *Neural Networks* p. 108241 (2025) [4](#), [10](#), [11](#), [12](#)
39. Mittal, A., Moorthy, A.K., Bovik, A.C.: No-reference image quality assessment in the spatial domain. *IEEE Trans. Image Process.* **21**(12), 4695–4708 (2012) [4](#), [6](#), [10](#)
40. Rafailov, R., Sharma, A., Mitchell, E., Manning, C.D., Ermon, S., Finn, C.: Direct preference optimization: Your language model is secretly a reward model. *Advances in neural information processing systems* **36**, 53728–53741 (2023) [3](#), [4](#), [7](#)
41. Su, S., Yan, Q., Zhu, Y., Zhang, C., Ge, X., Sun, J., Zhang, Y.: Blindly assess image quality in the wild guided by a self-adaptive hyper network. In: CVPR. pp. 3667–3676 (2020) [10](#)
42. Sun, S., Ren, W., Gao, X., Wang, R., Cao, X.: Restoring images in adverse weather conditions via histogram transformer. In: European Conference on Computer Vision. pp. 111–129. Springer (2024) [10](#), [11](#), [12](#)

43. Tarvainen, A., Valpola, H.: Mean teachers are better role models: Weight-averaged consistency targets improve semi-supervised deep learning results. In: *Advances in Neural Information Processing Systems*. pp. 1195–1204 (2017) [4](#)
44. Wang, J., Chan, K.C., Loy, C.C.: Exploring clip for assessing the look and feel of images. In: *Proceedings of the AAAI conference on artificial intelligence*. vol. 37, pp. 2555–2563 (2023) [4](#), [6](#), [10](#)
45. Wang, Y., Wang, H., Shen, Y., Fei, J., Li, W., Jin, G., Wu, L., Zhao, R., Le, X.: Semi-supervised semantic segmentation using unreliable pseudo-labels. In: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. pp. 4248–4257 (2022) [4](#)
46. Wang, Z., Bao, J., Gu, S., Chen, D., Zhou, W., Li, H.: Designdiffusion: High-quality text-to-design image generation with diffusion models. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. pp. 20906–20915 (June 2025) [4](#)
47. Wen, J., Cui, J., Zhao, Z., Yan, R., Gao, Z., Dou, L., Chen, B.M.: Syreanet: A physically guided underwater image enhancement framework integrating synthetic and real images. *arXiv preprint arXiv:2302.08269* (2023) [11](#), [12](#)
48. Wu, H., Zhang, Z., Zhang, W., Chen, C., Liao, L., Li, C., Gao, Y., Wang, A., Zhang, E., Sun, W., Yan, Q., Min, X., Zhai, G., Lin, W.: Q-align: Teaching llms for visual scoring via discrete text-defined levels. In: *ICML*. pp. 54015–54029 (2024) [10](#)
49. Wu, Q., Yang, J., Sun, K., Zhang, C., Zhang, Y., Salzmann, M.: Mixcycle: Mixup assisted semi-supervised 3d single object tracking with cycle consistency. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision*. pp. 13956–13966 (2023) [4](#)
50. Wu, Z., Xuan, H., Sun, C., Guan, W., Zhang, K., Yan, Y.: Semi-supervised video inpainting with cycle consistency constraints. In: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. pp. 22586–22595 (2023) [4](#)
51. Xiao, F., Hu, P., Xu, L., Guo, X., Qin, G., Shen, Y., Fang, C., Zhang, R., He, C., Farsiu, S.: Beyond ground-truth: Leveraging image quality priors for real-world image restoration. *CVPR* (2026) [4](#)
52. Xiao, F., Hu, S., Shen, Y., Fang, C., Huang, J., Tang, L., Yang, Z., Li, X., He, C.: A survey of camouflaged object detection and beyond. *CAAI AIR* **3** (2024) [4](#)
53. Xiao, T., Yuan, Y., Zhu, H., Li, M., Honavar, V.G.: Cal-dpo: Calibrated direct preference optimization for language model alignment. *Advances in Neural Information Processing Systems* **37**, 114289–114320 (2024) [4](#)
54. Xu, J., Wu, M., Hu, X., Fu, C.W., Dou, Q., Heng, P.A.: Towards real-world adverse weather image restoration: Enhancing clearness and semantics with vision-language models. In: *European Conference on Computer Vision*. pp. 147–164. Springer (2024) [4](#)
55. Xu, Z., Wang, Y., Lu, D., Luo, X., Yan, J., Zheng, Y., Tong, R.K.y.: Ambiguity-selective consistency regularization for mean-teacher semi-supervised medical image segmentation. *Medical Image Analysis* **88**, 102880 (2023) [4](#)
56. Yan, Q., Feng, Y., Zhang, C., Pang, G., Shi, K., Wu, P., Dong, W., Sun, J., Zhang, Y.: Hvi: A new color space for low-light image enhancement. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. pp. 5678–5687 (June 2025) [11](#), [12](#), [14](#)
57. Yang, F., Wu, K., Zhang, S., Jiang, G., Liu, Y., Zheng, F., Zhang, W., Wang, C., Zeng, L.: Class-aware contrastive semi-supervised learning. In: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. pp. 14421–14430 (2022) [4](#)

58. Yang, S., Wu, T., Shi, S., Lao, S., Gong, Y., Cao, M., Wang, J., Yang, Y.: Maniqa: Multi-dimension attention network for no-reference image quality assessment. In: CVPRW. pp. 1191–1200 (2022) [10](#), [13](#)
59. Yang, W., Wang, S., Fang, Y., Wang, Y., Liu, J.: From fidelity to perceptual quality: A semi-supervised approach for low-light image enhancement. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 3063–3072 (2020) [11](#), [12](#), [14](#)
60. Yu, Y., Chen, F., Yu, J., Kan, Z.: Lmt-gp: Combined latent mean-teacher and gaussian process for semi-supervised low-light image enhancement. In: European Conference on Computer Vision. pp. 261–279. Springer (2024) [11](#), [12](#), [14](#)
61. Zhang, F., Rangrej, S.B., Aumentado-Armstrong, T., Fazly, A., Levinshstein, A.: Augmenting perceptual super-resolution via image quality predictors. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 2311–2322 (2025) [4](#)
62. Zhang, W., Zhai, G., Wei, Y., Yang, X., Ma, K.: Blind image quality assessment via vision-language correspondence: A multitask learning perspective. In: CVPR. pp. 14071–14081 (2023) [10](#)
63. Zhu, H., Wu, H., Li, Y., Zhang, Z., Chen, B., Zhu, L., Fang, Y., Zhai, G., Lin, W., Wang, S.: Adaptive image quality assessment via teaching large multimodal model to compare. In: NeurIPS. pp. 32611–32629 (2024) [10](#)
64. Zhu, Y., Wang, T., Fu, X., Yang, X., Guo, X., Dai, J., Qiao, Y., Hu, X.: Learning weather-general and weather-specific features for image restoration under multiple adverse weather conditions. In: Proc. IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) (2023) [10](#), [12](#), [13](#)