

Cast and Attached Shadow Detection via Iterative Light and Geometry Reasoning

Shilin Hu¹, Jingyi Xu¹, Sagnik Das¹, Dimitris Samaras^{1†}, and Hieu Le^{2†}

¹ Stony Brook University, Stony Brook NY 11794, USA

² UNC Charlotte, Charlotte NC 28223, USA

{shilhu, jingyixu, sadas, samaras}@cs.stonybrook.edu, hle40@charlotte.edu

Abstract. Shadows encode rich information about scene geometry and illumination, yet existing methods either predict a unified shadow mask or overlook attached shadows entirely. We address this gap by proposing a framework for jointly detecting cast and attached shadows through explicit physical modeling of light direction and surface geometry under a dominant directional-light setting. Our approach is grounded in a simple observation: surfaces facing away from the light source tend to fall into shadow. We exploit the reciprocal relationship between shadow formation and light estimation to construct a closed feedback loop, a dual-module architecture in which a shadow detection module and a light estimation module iteratively refine each other. At each pass, updated light estimates, together with surface normals, produce partial attached shadow maps that guide detection, while improved shadow predictions sharpen light estimation. To support training and evaluation, we introduce a dataset of 1,458 images with manually annotated cast and attached shadow masks sourced from three existing benchmarks. Experiments demonstrate that our proposed method outperforms prior methods, with at least a 33% reduction in attached-shadow BER, while maintaining strong full-shadow and cast-shadow performance. Project page: https://shilin21.github.io/attached_detection/

Keywords: Shadow detection · Attached shadow

1 Introduction

Shadows define how we perceive shape, depth, and illumination in visual scenes. They are commonly categorized into two types: cast shadows, which appear on external surfaces and reveal spatial relations between objects [19, 31, 35, 43], and attached shadows, which form directly on an object’s surface when light is blocked due to the surface’s geometry and orientation relative to the light source [17, 71]. Together, these two shadow types encode complementary cues about scene structure and illumination, enriching visual realism and 3D perception [5, 52, 63, 65].

[†] Equal advising.

Table 1: Comparison of recent shadow detection datasets. Most existing datasets omit attached shadows or annotate them only implicitly, while our dataset provides explicit cast/attached shadow annotations as separate classes.

Dataset	Att.	#Imgs	Split
SBU [46]	✓ (occ.)	4.1K	✗
ISTD [47]	✗	1.9K	✗
SOBA [50]	✗	1.0K	✗
ViSha [3]	✗	11.7K	✗
SILT [61]	✓	0.6K	✗
CUHK [14]	✓	10.5K	✗
Ours	✓	1.5K	✓



Fig. 1: Effect of shadow mask granularity on removal quality. Using a unified shadow mask, ShadowFormer [10] over-attenuates object regions. Using separate cast and attached shadow masks sequentially enables targeted correction and recovers natural object color.

However, most modern shadow detection methods predict a unified shadow mask [14, 46, 61], or ignore attached shadows entirely [3, 47, 50] (see Tab. 1). As a result, **explicitly** modeling attached shadows remains largely unexplored in the literature, despite their potential usefulness in downstream applications. In 3D shape reconstruction, attached shadows provide strong cues about local surface orientation and therefore reveal valuable information about the underlying geometry of objects. In shadow removal, separate cast and attached masks enable more targeted correction of receiving surfaces and object regions. As a simple illustration, applying an existing method [10] with separate cast and attached masks improves removal quality over using a single unified mask (see Fig. 1).

Yet attached shadow detection is far more than a simple extension of existing shadow detection approaches. As we will show experimentally, naive adaptations of current methods [4, 14, 61, 68–70] fall short; see Fig. 2 for typical predictions. This difficulty arises because cast and attached shadows are governed by different formation mechanisms: cast shadows are often driven by object boundaries and produce relatively contiguous regions of contrast, amenable to appearance-based detection. Attached shadows, by contrast, arise from the interplay between surface orientation and illumination. They follow object geometry closely, often appearing fragmented with weak contrast, and resist detection without explicit reasoning about scene geometry and lighting.

We therefore introduce a framework that jointly detects cast and attached shadows by modeling their physical dependence on light direction and surface geometry. The key insight is simple: under direct illumination, surfaces facing away from the light receive no direct light and therefore form attached shadows. This relationship works in both directions. Given a light direction¹ and surface normals, we can estimate a partial attached-shadow map. Conversely, given

¹ We assume a single dominant directional light source.

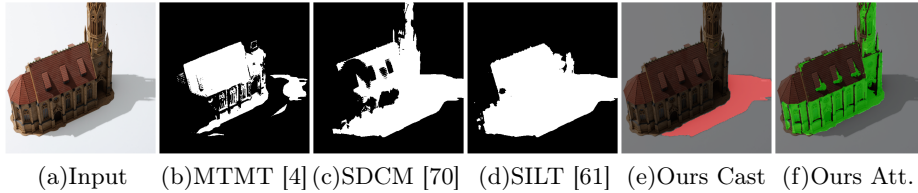


Fig. 2: Existing shadow detection methods primarily focus on cast shadow detection and do not differentiate between cast and attached shadow types (b, c, d). In contrast, our method accurately segments cast and attached shadows separately (e, f).

observed shadows and scene geometry, we can partially reason about the light direction. We exploit this reciprocal relationship through a closed-loop, dual-module architecture, in which a shadow detection module and a light estimation module mutually refine each other across iterative passes.

To support training and evaluation, we contribute a dataset of 1,458 images with manually annotated cast and attached shadow masks, drawn from three existing benchmarks [14, 44, 50]. Experiments confirm that our iterative, geometry- and light-aware formulation outperforms prior methods on both shadow types.

Our contributions are:

- A joint framework for detecting cast and attached shadows through shared learning with light estimation.
- An iterative refinement scheme in which lighting and shadow cues mutually reinforce each other.
- A curated dataset with separate annotations for cast and attached shadows, enabling standardized evaluation.
- Our method outperforms previous methods on both shadow types, with at least a 33% reduction in attached-shadow BER.

2 Related Work

Shadow detection and removal [6, 11–13, 16, 21–23, 28, 38] are fundamental tasks in computer vision, where accurate shadow estimation is crucial for tasks such as scene understanding [20, 24, 26, 27] and generation [54, 56–60]. Earlier shadow detection studies [30–34, 39] extensively explored the relationship between lighting and shadows. Panagopoulos *et al.* [35] were among the first to jointly model shadow detection and light estimation using coarse scene geometry. However, their approach relied on simplified geometry and idealized physical models, which limited its ability to handle the complexities and variations in real-world images.

Recent deep learning advancements have significantly improved shadow detection [2, 14, 15, 25, 47, 67, 68]. Ding *et al.* [8] introduced an attentive recurrent GAN that iteratively detects and removes shadows. Chen *et al.* [4] leveraged multiple cues from shadow regions, edges, and shadow counts, using a student-teacher model. Zhu *et al.* [69] noted that deep shadow detectors were overly

reliant on intensity cues and proposed a feature decomposition and reweighting approach to reduce intensity bias. SDCM [70] used a complementary mechanism to jointly detect shadow and non-shadow regions by transferring deactivated intermediate features between branches. Sun *et al.* [42] focused on preserving multiscale contrast information by mapping RAW images to sRGB representations of varying intensities. More recently, SwinShadow [53] leveraged a transformer-based architecture with a shifted window mechanism to better detect adjacent shadows by capturing hierarchical contextual details. In addition, several approaches have targeted more specialized tasks. [48–50] focused on the instance shadow detection task, concentrating on cast shadows associated with annotated objects. Wang *et al.* [51] proposed MetaShadow, which studies object-centered shadow reasoning by jointly addressing shadow detection, removal, and synthesis within a unified framework.

While these recent methods have significantly advanced shadow detection, they do not explicitly address the unique challenge of detecting attached shadows. Several methods and datasets acknowledge the importance of attached shadows. Yang *et al.* [61] annotated attached shadows for the SBU-test set and proposed a method that uses predicted shadows as updated training labels. Hu *et al.* [14] introduced a benchmark in which a large portion of the shadows are attached shadows. However, these approaches still treat cast and attached shadows uniformly, ignoring their distinctions.

Estimated lighting has been widely used in computer vision tasks. Existing illumination representations can be broadly categorized into three types: environment maps [40, 45], light probes [1, 37], and parameterized lighting models [7, 9, 64]. In this paper, we adopt a three-dimensional light-direction parameterization to model a dominant distant directional light, typical of outdoor daylight, and to provide a controlled, interpretable link between illumination and scene geometry. Importantly, shadow understanding under this single-direction assumption remains actively studied [55], and attached shadow detection is still challenging even in this setting.

3 Method

3.1 Framework Overview

Our framework detects both cast and attached shadows by jointly learning shadow detection and light estimation from a single RGB image and its corresponding normal map (Fig. 3). We define attached shadows as shadows on the object itself (typically on surface regions oriented away from the light), and cast shadows as shadows projected onto external surfaces. Attached shadows depend on surface orientation and lighting, motivating our geometry- and light-aware design. We assume a single scene-wide directional light shared across the image, and define ℓ to point from the light toward the scene.

Our shadow detection module is formulated as a multi-class segmentation task over {bg, cast, attached}, taking the RGB image and normal map as input.

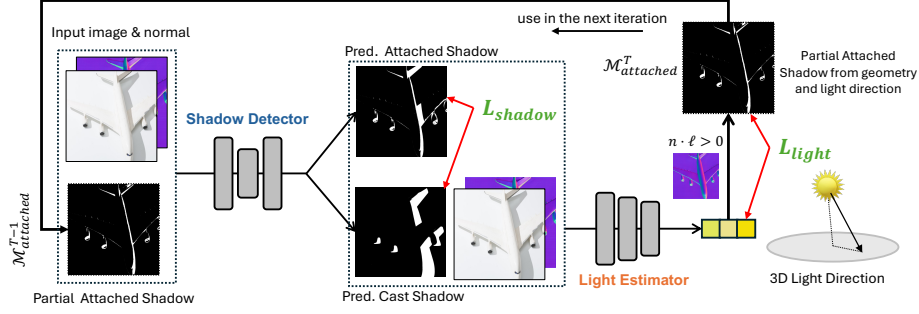


Fig. 3: Framework overview. Our network consists of a shadow detection module and a light estimation module. The detector predicts cast and attached shadows, while the light module estimates a 3D light direction and generates a geometry-based partial attached shadow map. This partial map is fed back to the detector as an additional input during iterative training, enabling progressive refinement.

Our light estimation module leverages the normal map and predicted shadows to estimate the light direction as a three-dimensional vector. Using the estimated light direction and surface normals, we generate a geometry-derived *partial attached shadow map* from a simple normal-light alignment cue. This partial map is incorporated as an additional input to refine shadow detection, while the predicted shadows provide auxiliary cues for light estimation. The two modules operate in a closed loop and are iteratively refined over multiple forward passes, improving both shadow detection and light estimation.

3.2 Iterative Training

Shadow and illumination serve as mutual cues for estimating each other [35]. Similarly, our framework jointly and iteratively refines shadow detection and light estimation.

Each training step consists of multiple iterations. In each iteration, we use the predicted light direction and surface normals to generate a *partial attached shadow map* derived from geometry and light direction. This map highlights pixels that are likely to belong to attached shadows based purely on local surface orientation and the light direction. It is not a complete attached shadow mask because it does not model visibility or geometric blocking; thus, it cannot capture cases where one surface region blocks light from reaching another. Still, it provides reliable local cues near contact regions and boundaries, and is fed to the shadow detection module in the next iteration as an additional input to predict the complete cast and attached shadow masks. Because this map is generated from the current light estimate during training, the detector learns to use imperfect geometry cues.

Specifically, given a light direction vector $\ell \in \mathbb{R}^3$ pointing from the light source toward the surface and surface normals $\mathbf{n} \in \mathbb{R}^{H \times W \times 3}$, the partial attached

shadow map, denoted as $\mathcal{M}_{\text{attached}}$, includes surface points oriented away from the light source as:

$$\mathcal{M}_{\text{attached}} = \begin{cases} 1, & \mathbf{n} \cdot \boldsymbol{\ell} > 0 \\ 0, & \mathbf{n} \cdot \boldsymbol{\ell} \leq 0 \end{cases} \quad (1)$$

This geometry-based map serves as a coarse attached-shadow prior. As the estimated light direction becomes more accurate across iterations, the map is progressively refined, leading to improved shadow detection and, in turn, better light estimation. In the first iteration, the partial attached-shadow map is set to an all-ones map.

3.3 Training Objectives

We jointly optimize the two modules with complementary objectives for shadow detection and light estimation. The shadow loss enforces accurate and well-separated segmentation of cast and attached shadows, while the light loss supervises the estimated light direction and regularizes it for physical plausibility. Together, these objectives encourage stable convergence of the iterative refinement process described above.

Shadow Detection Loss. We supervise the shadow detector at two complementary levels. First, for global shadow segmentation (shadow vs. background), we aggregate subclass logits with the log-sum-exp (LSE): $\text{LSE}(a, b) = \log(e^a + e^b)$. Let $z_{\text{bg}}, z_{\text{cast}}, z_{\text{att}} \in \mathbb{R}$ be per-pixel logits; we define:

$$s = \text{LSE}(z_{\text{cast}}, z_{\text{att}}) - z_{\text{bg}}. \quad (2)$$

We match s to the union ground-truth shadow mask $y_{\text{union}} \in \{0, 1\}$ using a binary cross-entropy (BCE) term plus a Dice [41] regularizer:

$$\mathcal{L}_{\text{seg}} = \text{BCE}(s, y_{\text{union}}) + \lambda_{\text{Dice}} \text{Dice}(s, y_{\text{union}}). \quad (3)$$

Second, to distinguish shadow types (cast vs. attached), we apply a standard multi-class cross-entropy (CE) on the full logits with per-pixel labels $y_{\text{type}} \in \{\text{bg}, \text{cast}, \text{att}\}$. In addition, we impose a class-conditional margin on the cast vs. attached logits to sharpen their separation. Let $d^{(i)} = z_{\text{cast}}^{(i)} - z_{\text{att}}^{(i)}$, and let Ω_{cast} and Ω_{att} be the index sets of pixels labeled cast and attached, respectively. With margin $m = 0.2$, we have:

$$\mathcal{L}_{\text{dist}} = \frac{1}{|\Omega_{\text{cast}}|} \sum_{i \in \Omega_{\text{cast}}} \max(0, m - d^{(i)}) + \frac{1}{|\Omega_{\text{att}}|} \sum_{i \in \Omega_{\text{att}}} \max(0, m + d^{(i)}), \quad (4)$$

$$\mathcal{L}_{\text{type}} = \text{CE}(\mathbf{z}, y_{\text{type}}) + \lambda_{\text{dist}} \mathcal{L}_{\text{dist}}. \quad (5)$$

Together, this supervision encourages robust recovery of overall shadow extent while explicitly teaching the model to separate attached from cast shadows. The total shadow loss is therefore:

$$\mathcal{L}_{\text{shadow}} = \mathcal{L}_{\text{seg}} + \mathcal{L}_{\text{type}}. \quad (6)$$

Light Estimation Loss. We constrain the light estimator with three complementary losses. First, we encourage light cues to localize attached shadows by aligning $\mathcal{M}_{\text{attached}}$ with the ground-truth attached shadow mask $y_{\text{att}} \in \{0, 1\}$ via a weighted BCE. Second, we regress the predicted light direction $\hat{\ell}$ to a heuristic target ℓ^* (from geometry; see Sec. 4 for details) with an ℓ_1 loss. Third, we enforce a unit norm:

$$\mathcal{L}_{\text{att}} = \text{BCE}_w(\mathcal{M}_{\text{attached}}, y_{\text{att}}), \quad \mathcal{L}_{\text{dir}} = \|\hat{\ell} - \ell^*\|_1, \quad \mathcal{L}_{\text{unit}} = (\|\hat{\ell}\|_2 - 1)^2. \quad (7)$$

The overall light loss is then:

$$\mathcal{L}_{\text{light}} = \lambda_{\text{att}}\mathcal{L}_{\text{att}} + \lambda_{\text{dir}}\mathcal{L}_{\text{dir}} + \lambda_{\text{unit}}\mathcal{L}_{\text{unit}}. \quad (8)$$

Finally, the full training objective is:

$$\mathcal{L}_{\text{total}} = \mathcal{L}_{\text{shadow}} + \mathcal{L}_{\text{light}}. \quad (9)$$

4 Dataset

We introduce the first dataset specifically curated for cast and attached shadow detection, consisting of 1,166 training and 292 test images. Each image is paired with (i) a surface normal map, (ii) a ray-based heuristic 3D light direction, (iii) binary cast and attached shadow masks, and (iv) an “undefined” shadow mask. Fig. 4 shows example annotations.

Our dataset is compiled from three public sources: 220 images from WSRD [44] (paired shadow/shadow-free images), 710 images from SOBA [50] (shadow images with instance-level cast masks), and 528 images from CUHK [14] (scene-level full-shadow masks).

Accurately annotating cast and attached shadows per pixel in natural scenes is highly labor-intensive due to mixed shadow types. To simplify annotation, we restrict fine-grained labeling to foreground objects: we manually mark the cast shadows and attached shadows associated with the annotated objects. All remaining shadow pixels not attributable to these objects are labeled as *undefined shadows*. We use undefined shadows only for full-shadow supervision and evaluation, and exclude them from cast/attached training and evaluation. This avoids ambiguous ownership and provides a fairer assessment of cast/attached separation. Additional details on data preparation are provided in the supplementary material.

Normal Maps. Since geometry is crucial for identifying attached shadows, we use a recent foundation model [62] to predict relative depth, which we convert to normal maps.

Object Masks. We extract object masks using BiRefNet [66]. These masks are used only for annotation and evaluation.

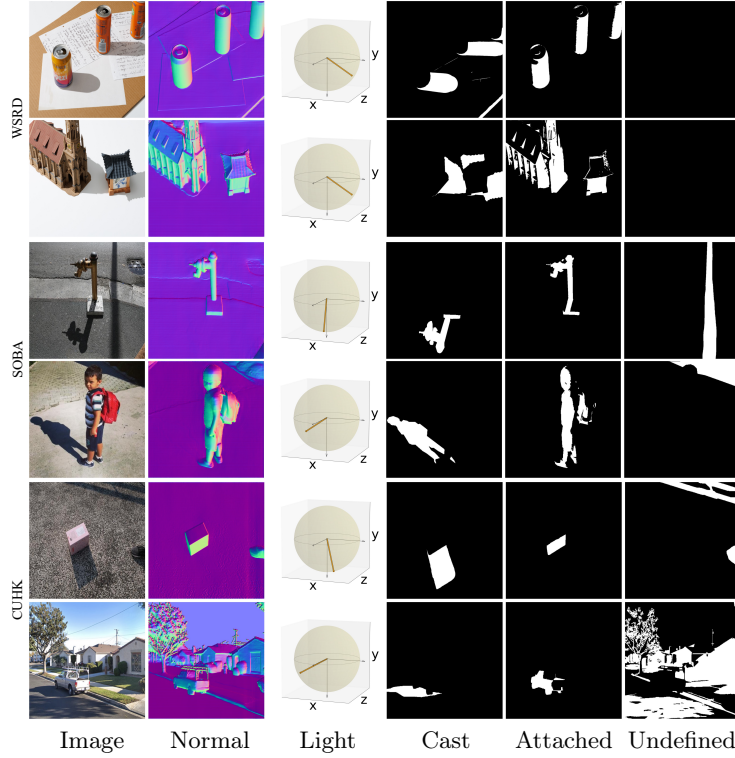


Fig. 4: Examples of our proposed dataset. We curate the dataset from WSRD [44], SOBA [50], and CUHK [14]. Each image is paired with its normal map, light direction, and cast/attached/“undefined” shadow masks. Light direction is camera-centric: $+x$ right, $+y$ down, $+z$ inward.

Cast and Attached Shadow Masks. Using the object masks, we manually annotate and refine cast and attached shadow masks for SOBA and CUHK. For WSRD, which lacks shadow masks, we derive full-shadow masks by color-space subtraction between shadow and shadow-free images, and then annotate cast and attached components.

Light Directions. WSRD images are captured under a calibrated fixed-lighting system [44]; therefore, we manually compute ground-truth 3D light directions. For SOBA and CUHK, we estimate a heuristic 3D direction in two steps. **(i) In-plane direction.** We compute a 2D image-plane direction by connecting the centroid of a foreground object to the centroid of its cast shadow. Under a directional-light assumption, a cast shadow can be viewed as the object translated along the illumination direction and projected onto the receiving surface; thus, the object-to-shadow displacement provides the illumination ray’s image-plane projection. **(ii) Depth sign.** We infer the third component from the relative depth between the object region and its cast shadow region. If the shadow lies

Table 2: Quantitative comparison with state-of-the-art methods on our dataset. We report BER↓ and F1↑ for full, cast, and attached shadows. Methods fine-tuned on our training set are marked with †. Best results are in **bold**.

Methods	Full		Cast		Attached	
	BER↓	F1↑	BER↓	F1↑	BER↓	F1↑
BDRAR [68]	21.29	70.53	7.23	79.64	35.41	52.32
DSD [67]	22.53	68.75	7.83	80.10	35.93	51.74
FSDNet [14]	24.46	66.09	10.51	79.16	37.53	46.24
MTMT [4]	33.10	50.00	15.16	74.66	42.02	38.10
FDRNet [69]	23.69	65.69	8.95	71.65	34.82	55.15
SDCM [70]	24.27	66.51	7.88	83.42	36.83	49.02
SILT [61]	20.20	71.51	4.64	80.39	32.70	60.23
BDRAR [†]	8.15	83.95	5.18	72.86	20.75	82.69
FSDNet [†]	10.17	85.10	5.82	81.94	19.49	80.57
FDRNet [†]	8.62	83.00	4.65	73.27	23.11	81.42
SILT [†]	6.51	82.98	4.52	68.31	26.57	80.72
Ours	6.50	86.61	5.04	85.46	12.92	86.49

on a surface deeper than the object, the illumination has a component pointing into the scene; if it lies shallower, the component points toward the camera. Combining this signed component with the in-plane vector and normalizing yields a 3D unit direction in our camera-centric frame (+ x right, + y down, + z inward).

5 Experiments and Results

Implementation Details. The proposed framework is implemented using PyTorch [36]. The shadow detection backbone follows SILT [61], and the light estimation module uses ConvNeXt-S [29]. We train with Adam [18] (learning rate 5×10^{-4} , 20 epochs, batch size 4). We set λ_{Dice} , λ_{dist} , λ_{att} , λ_{dir} , and λ_{unit} to 0.1, 0.2, 0.4, 0.5, and 0.1, respectively. Additional details are provided in the supplementary.

Evaluation Methods and Metrics. We compare against seven state-of-the-art methods: BDRAR [68], DSD [67], FSDNet [14], MTMT [4], FDRNet [69], SDCM [70], and SILT [61]. We evaluate all baselines with their SBU-pretrained weights and additionally fine-tune BDRAR, FSDNet, FDRNet, and SILT on our training set with full-shadow supervision (marked with † in Tab. 2).

We report the F1 score, $F1 = \frac{2TP}{2TP+FP+FN}$, and the balanced error rate (BER), $\text{BER} = 100 \times (1 - \frac{1}{2} [\frac{TP}{TP+FN} + \frac{TN}{TN+FP}])$, where TP , TN , FP , and FN denote true positives, true negatives, false positives, and false negatives.

For prior baselines that output a single shadow mask, we split their predictions using the object mask. Cast evaluation excludes the undefined region, and attached evaluation is restricted to the object region.

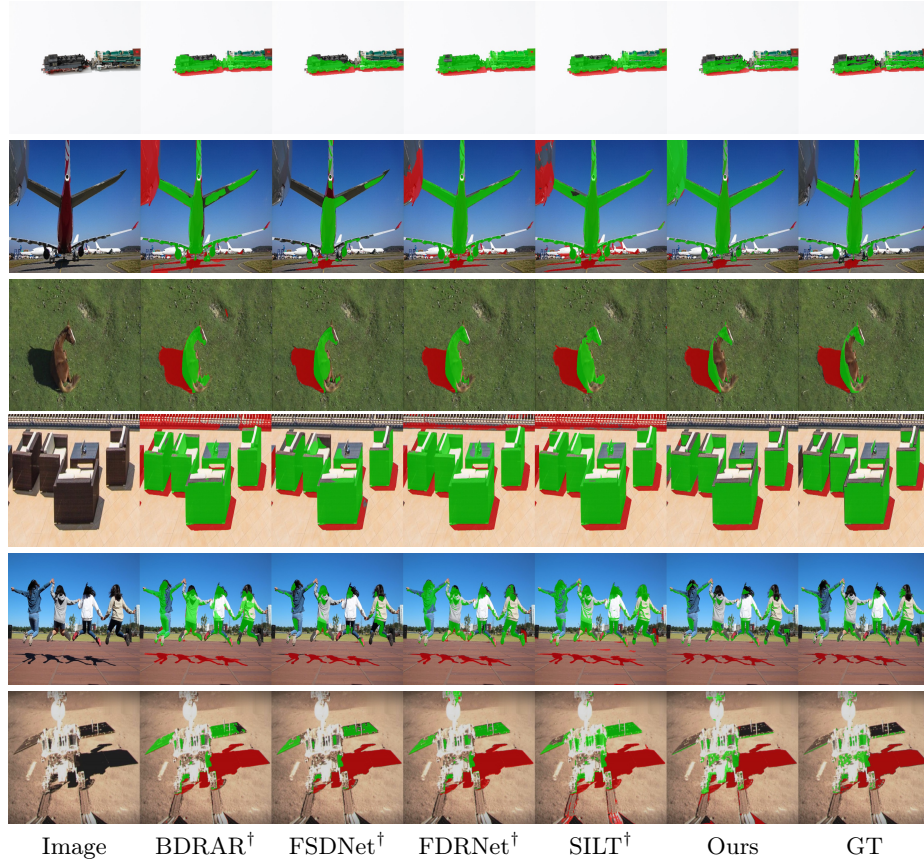


Fig. 5: Qualitative comparison of our method with retrained BDRAR [68], FSDNet [14], FDRNet [69], and SILT [61] on our dataset. Predicted masks are overlaid on the input image (*green: attached, red: cast*). Our geometry- and light-aware iterative framework produces more accurate attached-shadow boundaries.

5.1 Comparison with SOTA methods

Quantitative Comparisons. Tab. 2 reports BER and F1 scores for full, cast, and attached shadows. Our method achieves the best performance on full and attached shadows in both metrics and the strongest cast F1, while its cast BER remains comparable to prior approaches.

We first evaluate the SBU-pretrained baselines, which are cast-focused. All methods show a clear gap between cast and attached performance, with attached-shadow BER worse by about 20–30 points. These models primarily capture cast shadows, which often appear as contiguous, lower-intensity regions, and only detect attached shadows when they visually resemble cast shadows. This highlights the limitations of cast-biased training data for attached-shadow detection. We then retrain the baselines on our data using full-shadow supervision. This

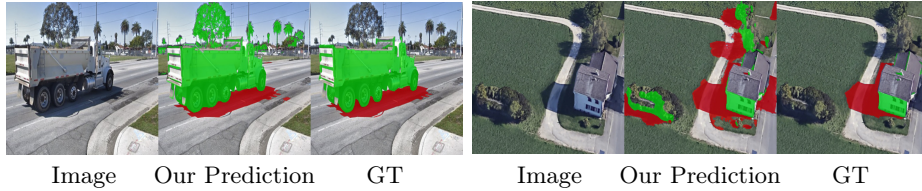


Fig. 6: Examples with foreground-only annotations. Although supervision is limited to the foreground objects, our method produces coherent cast and attached shadow predictions over the full image.

narrows the gap, but the baselines still struggle with attached shadows. A key limitation is that they treat the two shadow types uniformly and ignore that attached shadows are governed by surface orientation and contact rather than occlusion.

In contrast, our method leverages geometry and lighting by using surface normals and estimating a light direction. These physical cues let the model reason about where attached shadows should occur while preserving cast performance. As a result, attached shadow BER is reduced by at least 33% relative to retrained baselines, while maintaining strong performance on full and cast shadows.

Qualitative Comparisons. Fig. 5 shows visual comparisons between four retrained baselines and our method. The baselines produce a single shadow mask, so we use the object mask to split their outputs into cast and attached regions. Our method follows occluder geometry and the estimated light direction, producing geometry-consistent attached-shadow boundaries on surfaces facing away from the light (rows 1–2). In contrast, the baselines rely mainly on image appearance and tend to mark dark regions as shadows, even when those surfaces are lit (rows 3–5). By using surface normals, an estimated light direction, and separate predictions for cast and attached shadows, our method improves attached shadow quality while maintaining strong cast predictions.

We further show qualitative results in Fig. 6 where the ground-truth annotations cover only foreground objects, yet our method correctly assigns cast and attached labels across the remaining undefined shadow regions. This suggests a potential path toward reducing dense annotation effort through pseudo-labels or weak supervision.

5.2 Ablation Studies

To validate the effectiveness of our framework design, we conduct ablation studies to analyze the contributions of (i) surface-normal guidance, (ii) the iterative learning scheme, and (iii) each loss component.

Effect of geometry guidance. We examine the role of surface normals as geometry guidance. Removing surface normals reduces our model to a standard multi-class segmentation network trained with class-specific labels. This substantially

Table 3: Ablation of our geometry- and light-aware framework. We evaluate the impact of (i) removing surface normals, (ii) varying the number of refinement iterations, and (iii) disabling individual loss terms. We report $\text{BER}\downarrow$ and $\text{F1}\uparrow$ for full, cast, and attached shadow detection.

Setting	Full		Cast		Attached	
	$\text{BER}\downarrow$	$\text{F1}\uparrow$	$\text{BER}\downarrow$	$\text{F1}\uparrow$	$\text{BER}\downarrow$	$\text{F1}\uparrow$
Geometry						
No Normal	7.60	85.46	7.82	79.33	20.87	77.07
Iterations						
Non-iterative	7.57	85.99	6.14	87.75	15.10	83.91
2 iterations	6.63	87.77	5.41	86.75	13.22	85.63
3 iterations (final)	6.50	86.61	5.04	85.46	12.92	86.49
5 iterations	6.11	87.14	5.83	85.57	13.51	85.74
Loss terms						
No $\mathcal{L}_{\text{dist}}$	7.19	86.25	6.16	85.69	13.45	85.38
No $\mathcal{L}_{\text{att}}$	6.89	87.03	5.81	85.26	14.61	84.24
No $\mathcal{L}_{\text{dir}}$	7.50	86.42	4.97	85.02	15.26	83.59

degrades performance across metrics, with the largest drop on attached shadows, yielding results comparable to retrained baselines. This confirms that surface normals provide strong spatial priors for separating cast and attached shadows.

Effect of iterative learning. We evaluate a non-iterative variant that jointly trains shadow detection and light estimation with normal maps, but does not feed the predicted partial attached map back into the detector. As shown in Tab. 3, this setting outperforms the variant without normal guidance, confirming the benefit of geometry awareness. We then vary the number of refinement iterations (2, 3, and 5). Iterative feedback improves predictions: two iterations provide the largest gain, and performance saturates at three iterations, as also shown by the intermediate results in Fig. 7. Increasing to five iterations slightly improves full-shadow accuracy but reduces class-specific performance, suggesting over-refinement of the partial map. In terms of efficiency, runtime increases approximately linearly with the number of iterations; our 3-step setting is $3\times$ slower than the non-iterative model (0.55s vs. 0.18s per image). Overall, geometry provides the key prior, and iterative refinement further improves segmentation.

Effect of loss design. We ablate each loss term individually during training. Removing any single term reduces performance, indicating that the losses are complementary. $\mathcal{L}_{\text{dist}}$ encourages separation between cast and attached logits via a class-conditional margin, leading to a clearer decision boundary. The light-related losses, $\mathcal{L}_{\text{att}}$ and $\mathcal{L}_{\text{dir}}$, guide the light estimator to produce spatially coherent partial maps by aligning outputs with ground-truth cues and enforcing directional consistency. Together, these objectives promote effective geometry-light interaction and more reliable shadow predictions.

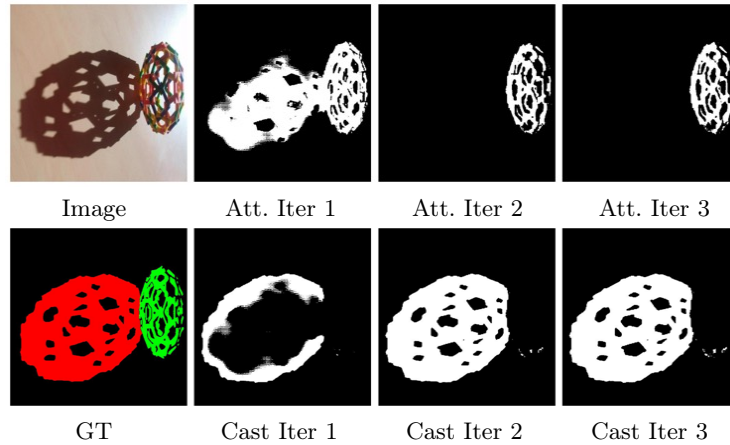


Fig. 7: Effectiveness of the iterative learning scheme. Our proposed feedback mechanism progressively refines the shadow predictions, showing clear visual improvements at the second iteration and stabilization at the third.

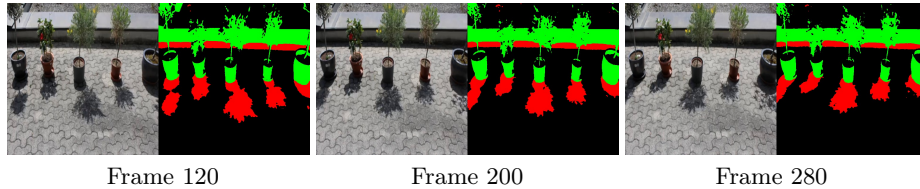


Fig. 8: Cross-dataset generalization on SBU-TimeLapse [23]. We apply our model to a static-camera video with time-varying shadows and visualize predictions at different time steps. The results remain temporally consistent and show coherent cast and attached shadow separation across the sequence.

5.3 Cross-dataset Generalization

To assess generalization beyond our training data, we conduct experiments on SBU-TimeLapse [23], a static-scene video dataset with time-varying shadows. We run our model on each frame independently. As shown in Fig. 8, our predictions remain temporally consistent while producing accurate cast and attached shadow separation across the sequence, indicating robustness to temporal shadow variation. Additional qualitative results are provided in the supplementary.

5.4 Shadow Removal Application

Shadow masks are widely used in downstream tasks. To demonstrate the practical value of separating attached shadows, we study shadow removal as an example. We guide an off-the-shelf shadow removal model [10] using a *joint* shadow mask, and then refine the result within the object region using the *attached* shadow mask. Because our data do not provide shadow-free ground truth for

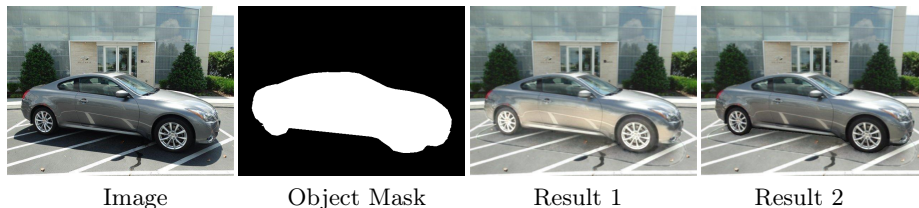


Fig. 9: User study example. Each question shows the input image and object mask, along with two candidate results: one guided by a joint shadow mask and the other guided by separate cast/attached masks with object-region refinement. Candidate order is randomized, and participants select the more natural output.

Table 4: User study on shadow removal. Participants were asked to choose the more natural result between the joint-mask guidance (Joint) and separate-mask guidance with attached-region refinement (Separate).

Method Preference (%)	
Joint	19.3
Separate	80.7

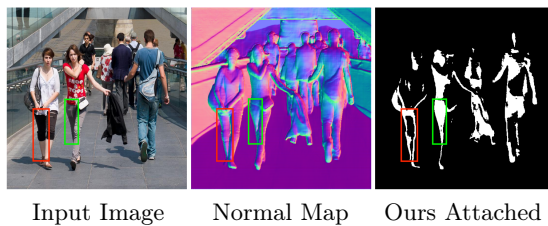


Fig. 10: Limitations. Attached-shadow errors can arise from inaccurate surface normals, such as flipped normal directions that lead to incorrect attached-shadow predictions.

quantitative evaluation, we conduct a user study in which participants choose the more natural of the two outputs (see Fig. 9). As shown in Tab. 4, participants prefer the result that uses the attached mask for refinement in 80.7% of comparisons. This outcome underscores the importance of separating cast and attached shadows: it enables targeted refinement on the object while removing cast shadows, leading to more realistic results.

6 Limitations

As an initial step toward detecting both cast and attached shadows, our framework has certain limitations. First, it relies on scene geometry derived from normal maps generated by an off-the-shelf model [62]. However, this foundation model does not always yield accurate surface orientations, and such errors can propagate to our predictions, as illustrated in Fig. 10. Improving the reliability of geometric cues through better normal estimation or geometry refinement is a promising direction for future work. Second, our approach assumes a single directional light source, which may not generalize well to complex illumination settings, including multi-source indoor lighting, soft shadows from extended light sources, and low-light/night-time scenes with strong indirect illumination. Future work could extend the lighting formulation beyond a single directional

source by modeling multiple or area lights, and broaden training data to cover more diverse illumination conditions.

7 Conclusion

We introduced a geometry- and light-aware framework for separately detecting cast and attached shadows, addressing the gap where most prior shadow detectors predict a unified mask or remain cast-focused. Motivated by the relationship between illumination, surface geometry, and shadows, our method jointly learns shadow segmentation and light estimation: the estimated light direction and surface normals yield a geometry-derived partial attached shadow map, which is fed back to refine detection, while predicted shadows provide cues for updating the light direction. This closed-loop, iterative formulation, together with explicit supervision for cast/attached separation, consistently improves attached-shadow detection while maintaining strong cast-shadow performance, as shown by comparisons to state-of-the-art baselines. To support future research and model development, we introduce a new dataset with high-quality separate annotations for cast and attached shadows.

Acknowledgements

This work was supported in part by the CCI startup fund at UNC Charlotte and NSF Grant IIS-2212046.

References

1. Bai, J., He, Z., Yang, S., Guo, J., Chen, Z., Zhang, Y., Guo, Y.: Local-to-global panorama inpainting for locale-aware indoor lighting prediction. *IEEE Trans. Vis. Comput. Graph.* **29**(11), 4405–4416 (2023). <https://doi.org/10.1109/TVCG.2023.3320233>
2. Chen, T., Zhu, L., Ding, C., Cao, R., Wang, Y., Zhang, S., Li, Z., Sun, L., Zang, Y., Mao, P.: SAM-adapter: Adapting segment anything in underperformed scenes. In: *Proc. IEEE/CVF Int. Conf. Comput. Vis. Workshops (ICCVW)*. pp. 3359–3367 (2023). <https://doi.org/10.1109/ICCVW60793.2023.00361>
3. Chen, Z., Wan, L., Zhu, L., Shen, J., Fu, H., Liu, W., Qin, J.: Triple-cooperative video shadow detection. In: *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR)*. pp. 2714–2723 (2021). <https://doi.org/10.1109/CVPR46437.2021.00274>
4. Chen, Z., Zhu, L., Wan, L., Wang, S., Feng, W., Heng, P.A.: A multi-task mean teacher for semi-supervised shadow detection. In: *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR)*. pp. 5610–5619 (2020). <https://doi.org/10.1109/CVPR42600.2020.00565>
5. Chuang, Y.Y., Goldman, D.B., Curless, B., Salesin, D.H., Szeliski, R.: Shadow matting and compositing. *ACM Trans. Graph.* **22**(3), 494–500 (2003). <https://doi.org/10.1145/1201775.882298>

6. Cun, X., Pun, C.M., Shi, C.: Towards ghost-free shadow removal via dual hierarchical aggregation network and shadow matting GAN. In: Proc. AAAI Conf. Artif. Intell. vol. 34, pp. 10680–10687 (2020). <https://doi.org/10.1609/aaai.v34i07.6695>
7. Dastjerdi, M.R.K., Eisenmann, J., Hold-Geoffroy, Y., Lalonde, J.F.: EverLight: Indoor-outdoor editable HDR lighting estimation. In: Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV). pp. 7386–7395 (2023). <https://doi.org/10.1109/ICCV51070.2023.00682>
8. Ding, B., Long, C., Zhang, L., Xiao, C.: ARGAN: Attentive recurrent generative adversarial network for shadow detection and removal. In: Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV). pp. 10212–10221 (2019). <https://doi.org/10.1109/ICCV.2019.01031>
9. Gardner, M.A., Hold-Geoffroy, Y., Sunkavalli, K., Gagné, C., Lalonde, J.F.: Deep parametric indoor lighting estimation. In: Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV). pp. 7174–7182 (2019). <https://doi.org/10.1109/ICCV.2019.00727>
10. Guo, L., Huang, S., Liu, D., Cheng, H., Wen, B.: ShadowFormer: Global context helps shadow removal. In: Proc. AAAI Conf. Artif. Intell. vol. 37, pp. 710–718 (2023). <https://doi.org/10.1609/aaai.v37i1.25148>
11. Guo, L., Wang, C., Yang, W., Huang, S., Wang, Y., Pfister, H., Wen, B.: ShadowDiffusion: When degradation prior meets diffusion model for shadow removal. In: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR). pp. 14049–14058 (2023). <https://doi.org/10.1109/CVPR52729.2023.01350>
12. Hu, S., Le, H., Athar, S., Das, S., Samaras, D.: Shadow removal refinement via material-consistent shadow edges. In: Proc. IEEE/CVF Winter Conf. Appl. Comput. Vis. (WACV). pp. 2631–2641 (2025). <https://doi.org/10.1109/WACV61041.2025.00261>
13. Hu, X., Jiang, Y., Fu, C.W., Heng, P.A.: Mask-ShadowGAN: Learning to remove shadows from unpaired data. In: Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV). pp. 2472–2481 (2019). <https://doi.org/10.1109/ICCV.2019.00256>
14. Hu, X., Wang, T., Fu, C.W., Jiang, Y., Wang, Q., Heng, P.A.: Revisiting shadow detection: A new benchmark dataset for complex world. IEEE Trans. Image Process. **30**, 1925–1934 (2021). <https://doi.org/10.1109/TIP.2021.3049331>
15. Hu, X., Zhu, L., Fu, C.W., Qin, J., Heng, P.A.: Direction-aware spatial context features for shadow detection. In: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR). pp. 7454–7462 (2018). <https://doi.org/10.1109/CVPR.2018.00778>
16. Jin, Y., Ye, W., Yang, W., Yuan, Y., Tan, R.T.: DeS3: Adaptive attention-driven self and soft shadow removal using ViT similarity. In: Proc. AAAI Conf. Artif. Intell. vol. 38, pp. 2634–2642 (2024). <https://doi.org/10.1609/aaai.v38i3.28041>
17. KaewTraKulPong, P., Bowden, R.: An improved adaptive background mixture model for real-time tracking with shadow detection. In: Video-Based Surveillance Systems, pp. 135–144. Springer (2002). https://doi.org/10.1007/978-1-4615-0913-4_11
18. Kingma, D.P., Ba, J.: Adam: A method for stochastic optimization (2017). <https://doi.org/10.48550/arXiv.1412.6980>
19. Lalonde, J.F., Efros, A.A., Narasimhan, S.G.: Detecting ground shadows in outdoor consumer photographs. In: Proc. Eur. Conf. Comput. Vis. (ECCV). pp. 322–335 (2010). https://doi.org/10.1007/978-3-642-15552-9_24

20. Le, H., Nguyen, V., Yu, C.P., Samaras, D.: Geodesic distance histogram feature for video segmentation. In: Proc. Asian Conf. Comput. Vis. (ACCV). pp. 275–290 (2016). https://doi.org/10.1007/978-3-319-54181-5_18
21. Le, H., Samaras, D.: Shadow removal via shadow image decomposition. In: Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV). pp. 8577–8586 (2019). <https://doi.org/10.1109/ICCV.2019.00867>
22. Le, H., Samaras, D.: From shadow segmentation to shadow removal. In: Proc. Eur. Conf. Comput. Vis. (ECCV). pp. 264–281 (2020). https://doi.org/10.1007/978-3-030-58621-8_16
23. Le, H., Samaras, D.: Physics-based shadow image decomposition for shadow removal. IEEE Trans. Pattern Anal. Mach. Intell. **44**(12), 9088–9101 (2022). <https://doi.org/10.1109/TPAMI.2021.3124934>
24. Le, H., Samaras, D., Lynch, H.J.: A convolutional neural network architecture designed for the automated survey of seabird colonies. Remote Sens. Ecol. Conserv. **8**(2), 251–262 (2022). <https://doi.org/10.1002/rse2.240>
25. Le, H., Vicente, T.F.Y., Nguyen, V., Hoai, M., Samaras, D.: A+D Net: Training a shadow detector with adversarial shadow attenuation. In: Proc. Eur. Conf. Comput. Vis. (ECCV). pp. 680–696 (2018). https://doi.org/10.1007/978-3-030-01216-8_41
26. Le, H., Yu, C.P., Zelinsky, G., Samaras, D.: Co-localization with category-consistent features and geodesic distance propagation. In: Proc. IEEE Int. Conf. Comput. Vis. Workshops (ICCVW). pp. 1103–1112 (2017). <https://doi.org/10.1109/ICCVW.2017.134>
27. Le, H.M., Goncalves, B., Samaras, D., Lynch, H.: Weakly labeling the antarctic: The penguin colony case. In: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. Workshops (CVPRW). pp. 18–25 (2019). <https://doi.org/10.48550/arXiv.1905.03313>
28. Li, X., Guo, Q., Abdelfattah, R., Lin, D., Feng, W., Tsang, I., Wang, S.: Leveraging inpainting for single-image shadow removal. In: Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV). pp. 13009–13018 (2023). <https://doi.org/10.1109/ICCV51070.2023.01200>
29. Liu, Z., Mao, H., Wu, C.Y., Feichtenhofer, C., Darrell, T., Xie, S.: A ConvNet for the 2020s. In: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR). pp. 11966–11976 (2022). <https://doi.org/10.1109/CVPR52688.2022.01167>
30. Okabe, T., Sato, I., Sato, Y.: Attached shadow coding: Estimating surface normals from shadows under unknown reflectance and lighting conditions. In: Proc. IEEE Int. Conf. Comput. Vis. (ICCV). pp. 1693–1700 (2009). <https://doi.org/10.1109/ICCV.2009.5459381>
31. Panagopoulos, A., Samaras, D., Paragios, N.: Robust shadow and illumination estimation using a mixture model. In: Proc. IEEE Conf. Comput. Vis. Pattern Recog. Workshops (CVPRW) (2009). <https://doi.org/10.1109/CVPRW.2009.5206665>
32. Panagopoulos, A., Vicente, T.F.Y., Samaras, D.: Illumination estimation from shadow borders. In: Proc. IEEE Int. Conf. Comput. Vis. Workshops (ICCVW). pp. 798–805 (2011). <https://doi.org/10.1109/ICCVW.2011.6130334>
33. Panagopoulos, A., Wang, C., Samaras, D., Paragios, N.: Illumination estimation and cast shadow detection through a higher-order graphical model. In: Proc. IEEE Conf. Comput. Vis. Pattern Recog. (CVPR). pp. 673–680 (2011). <https://doi.org/10.1109/CVPR.2011.5995585>

34. Panagopoulos, A., Wang, C., Samaras, D., Paragios, N.: Estimating shadows with the bright channel cue. In: Trends and Topics in Computer Vision. pp. 1–12 (2012). https://doi.org/10.1007/978-3-642-35740-4_1
35. Panagopoulos, A., Wang, C., Samaras, D., Paragios, N.: Simultaneous cast shadows, illumination and geometry inference using hypergraphs. *IEEE Trans. Pattern Anal. Mach. Intell.* **35**(2), 437–449 (2013). <https://doi.org/10.1109/TPAMI.2012.110>
36. Paszke, A., Gross, S., Massa, F., Lerer, A., Bradbury, J., Chanan, G., Killeen, T., Lin, Z., Gimelshein, N., Antiga, L., Desmaison, A., Köpf, A., Yang, E., DeVito, Z., Raison, M., Tejani, A., Chilamkurthy, S., Steiner, B., Fang, L., Bai, J., Chintala, S.: PyTorch: An imperative style, high-performance deep learning library. In: Adv. Neural Inf. Process. Syst. (NeurIPS). vol. 32 (2019). <https://doi.org/10.48550/arXiv.1912.01703>
37. Phongthawee, P., Chinchuthakun, W., Sinsunthithet, N., Jampani, V., Raj, A., Khungurn, P., Suwajanakorn, S.: DiffusionLight: Light probes for free by painting a chrome ball. In: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR). pp. 98–108 (2024). <https://doi.org/10.1109/CVPR52733.2024.00018>
38. Qu, L., Tian, J., He, S., Tang, Y., Lau, R.W.H.: DeshadowNet: A multi-context embedding deep network for shadow removal. In: Proc. IEEE Conf. Comput. Vis. Pattern Recog. (CVPR). pp. 2308–2316 (2017). <https://doi.org/10.1109/CVPR.2017.248>
39. Sato, I., Sato, Y., Ikeuchi, K.: Illumination from shadows. *IEEE Trans. Pattern Anal. Mach. Intell.* **25**(3), 290–300 (2003). <https://doi.org/10.1109/TPAMI.2003.1182093>
40. Somanath, G., Kurz, D.: HDR environment map estimation for real-time augmented reality. In: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR). pp. 11293–11301 (2021). <https://doi.org/10.1109/CVPR46437.2021.01114>
41. Sudre, C.H., Li, W., Vercauteren, T., Ourselin, S., Cardoso, M.J.: Generalised dice overlap as a deep learning loss function for highly unbalanced segmentations. In: Proc. Deep Learn. Med. Image Anal. Multimodal Learn. Clin. Decis. Support. pp. 240–248 (2017). https://doi.org/10.1007/978-3-319-67558-9_28
42. Sun, J., Xu, K., Pang, Y., Zhang, L., Lu, H., Hancke, G., Lau, R.: Adaptive illumination mapping for shadow detection in raw images. In: Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV). pp. 12663–12672 (2023). <https://doi.org/10.1109/ICCV51070.2023.01167>
43. Sunkavalli, K., Zickler, T., Pfister, H.: Visibility subspaces: Uncalibrated photometric stereo with shadows. In: Proc. Eur. Conf. Comput. Vis. (ECCV). pp. 251–264 (2010). https://doi.org/10.1007/978-3-642-15552-9_19
44. Vasluianu, F.A., Seizinger, T., Zhou, Z., Wu, Z., Chen, C., Timofte, R., et al.: NTIRE 2024 image shadow removal challenge report. In: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. Workshops (CVPRW). pp. 6547–6570 (2024). <https://doi.org/10.1109/CVPRW63382.2024.00654>
45. Verbin, D., Mildenhall, B., Hedman, P., Barron, J.T., Zickler, T., Srinivasan, P.P.: Eclipse: Disambiguating illumination and materials using unintended shadows. In: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR). pp. 77–86 (2024). <https://doi.org/10.1109/CVPR52733.2024.00016>
46. Vicente, T.F.Y., Hou, L., Yu, C.P., Hoai, M., Samaras, D.: Large-scale training of shadow detectors with noisily-annotated shadow examples. In: Proc. Eur. Conf. Comput. Vis. (ECCV). pp. 816–832 (2016). https://doi.org/10.1007/978-3-319-46466-4_49

47. Wang, J., Li, X., Yang, J.: Stacked conditional generative adversarial networks for jointly learning shadow detection and shadow removal. In: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR). pp. 1788–1797 (2018). <https://doi.org/10.1109/CVPR.2018.00192>
48. Wang, T., Hu, X., Fu, C.W., Heng, P.A.: Single-stage instance shadow detection with bidirectional relation learning. In: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR). pp. 1–11 (2021). <https://doi.org/10.1109/CVPR46437.2021.00007>
49. Wang, T., Hu, X., Heng, P.A., Fu, C.W.: Instance shadow detection with a single-stage detector. *IEEE Trans. Pattern Anal. Mach. Intell.* **45**(3), 3259–3273 (2023). <https://doi.org/10.1109/TPAMI.2022.3185628>
50. Wang, T., Hu, X., Wang, Q., Heng, P.A., Fu, C.W.: Instance shadow detection. In: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR). pp. 1877–1886 (2020). <https://doi.org/10.1109/CVPR42600.2020.00195>
51. Wang, T., Zhang, J., Zheng, H., Ding, Z., Cohen, S., Lin, Z., Xiong, W., Fu, C.W., Figueroa, L., Kim, S.Y.: MetaShadow: Object-centered shadow detection, removal, and synthesis. In: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR). pp. 28252–28262 (2025). <https://doi.org/10.1109/CVPR52734.2025.02631>
52. Wang, Y., Curless, B.L., Seitz, S.M.: People as scene probes. In: Proc. Eur. Conf. Comput. Vis. (ECCV). pp. 438–454 (2020). https://doi.org/10.1007/978-3-030-58607-2_26
53. Wang, Y., Liu, S., Li, L., Zhou, W., Li, H.: SwinShadow: Shifted window for ambiguous adjacent shadow detection. *ACM Trans. Multimedia Comput. Commun. Appl.* **20**(11), 1–20 (2024). <https://doi.org/10.1145/3688803>
54. Wu, H., Xu, J., Le, H., Samaras, D.: Importance-based token merging for efficient image and video generation. In: Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV). pp. 4983–4995 (2025). <https://doi.org/10.1109/ICCV51701.2025.00474>
55. Xing, Z., Wang, T., Hu, X., Wu, H., Fu, C.W., Heng, P.A.: Video instance shadow detection under the sun and sky. *IEEE Trans. Image Process.* **33**, 5715–5726 (2024). <https://doi.org/10.1109/TIP.2024.3468877>
56. Xu, J., Le, H.: Generating representative samples for few-shot classification. In: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR). pp. 8993–9003 (2022). <https://doi.org/10.1109/CVPR52688.2022.00880>
57. Xu, J., Le, H., Huang, M., Athar, S., Samaras, D.: Variational feature disentangling for fine-grained few-shot classification. In: Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV). pp. 8792–8801 (2021). <https://doi.org/10.1109/ICCV48922.2021.00869>
58. Xu, J., Le, H., Nguyen, V., Ranjan, V., Samaras, D.: Zero-shot object counting. In: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR). pp. 15548–15557 (2023). <https://doi.org/10.1109/CVPR52729.2023.01492>
59. Xu, J., Le, H., Samaras, D.: Generating features with increased crop-related diversity for few-shot object detection. In: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR). pp. 19713–19722 (2023). <https://doi.org/10.1109/CVPR52729.2023.01888>
60. Xu, J., Le, H., Samaras, D.: Assessing sample quality via the latent space of generative models. In: Proc. Eur. Conf. Comput. Vis. (ECCV). pp. 449–464 (2024). https://doi.org/10.1007/978-3-031-73202-7_26
61. Yang, H., Wang, T., Hu, X., Fu, C.W.: SILT: Shadow-aware iterative label tuning for learning to detect shadows from noisy labels. In: Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV). pp. 12641–12652 (2023). <https://doi.org/10.1109/ICCV51070.2023.01165>

62. Yang, L., Kang, B., Huang, Z., Zhao, Z., Xu, X., Feng, J., Zhao, H.: Depth anything V2. In: *Adv. Neural Inf. Process. Syst. (NeurIPS)*. vol. 37, pp. 21875–21911 (2024). <https://doi.org/10.52202/079017-0688>
63. Yang, Y., Hu, S., Wu, H., Baldrich, R., Samaras, D., Vanrell, M.: mli-NeRF: Multi-light intrinsic-aware neural radiance fields. In: *Proc. Int. Conf. 3D Vision (3DV)*. pp. 587–596 (2025). <https://doi.org/10.1109/3DV66043.2025.00059>
64. Zhang, J., Sunkavalli, K., Hold-Geoffroy, Y., Hadap, S., Eisenman, J., Lalonde, J.F.: All-weather deep outdoor lighting estimation. In: *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR)*. pp. 10150–10158 (2019). <https://doi.org/10.1109/CVPR.2019.01040>
65. Zhang, W., Zhao, X., Morvan, J.M., Chen, L.: Improving shadow suppression for illumination robust face recognition. *IEEE Trans. Pattern Anal. Mach. Intell.* **41**(3), 611–624 (2019). <https://doi.org/10.1109/TPAMI.2018.2803179>
66. Zheng, P., Gao, D., Fan, D.P., Liu, L., Laaksonen, J., Ouyang, W., Sebe, N.: Bilateral reference for high-resolution dichotomous image segmentation. *CAAI Artif. Intell. Res.* p. 9150038 (2024). <https://doi.org/10.26599/AIR.2024.9150038>
67. Zheng, Q., Qiao, X., Cao, Y., Lau, R.W.H.: Distraction-aware shadow detection. In: *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recog. (CVPR)*. pp. 5162–5171 (2019). <https://doi.org/10.1109/CVPR.2019.00531>
68. Zhu, L., Deng, Z., Hu, X., Fu, C.W., Xu, X., Qin, J., Heng, P.A.: Bidirectional feature pyramid network with recurrent attention residual modules for shadow detection. In: *Proc. Eur. Conf. Comput. Vis. (ECCV)*. pp. 122–137 (2018). https://doi.org/10.1007/978-3-030-01231-1_8
69. Zhu, L., Xu, K., Ke, Z., Lau, R.W.H.: Mitigating intensity bias in shadow detection via feature decomposition and reweighting. In: *Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV)*. pp. 4682–4691 (2021). <https://doi.org/10.1109/ICCV48922.2021.00466>
70. Zhu, Y., Fu, X., Cao, C., Wang, X., Sun, Q., Zha, Z.J.: Single image shadow detection via complementary mechanism. In: *Proc. ACM Int. Conf. Multimedia (ACM MM)*. pp. 6717–6726 (2022). <https://doi.org/10.1145/3503161.3547904>
71. Zhu, Z., Woodcock, C.E.: Object-based cloud and cloud shadow detection in Landsat imagery. *Remote Sens. Environ.* **118**, 83–94 (2012). <https://doi.org/10.1016/j.rse.2011.10.028>