

Unsupervised Deep Video Denoising

Dev Yashpal Sheth^{1*}, Sreyas Mohan^{2*}, Joshua L. Vincent³, Ramon Manzorro³, Peter A. Crozier³,
 Mitesh M. Khapra^{1,6}, Eero P. Simoncelli^{2,4,5}, Carlos Fernandez-Granda^{2,5}

¹Indian Institute of Technology Madras, India, ²Center for Data Science, New York University,

³School for Engineering of Matter, Transport, and Energy, ASU,

⁴Center for Neural Science, NYU and Flatiron Institute, Simons Foundation,

⁵Courant Institute of Mathematical Sciences, NYU, ⁶Robert Bosch Center for Data Science and AI.

Abstract

Deep convolutional neural networks (CNNs) for video denoising are typically trained with supervision, assuming the availability of clean videos. However, in many applications, such as microscopy, noiseless videos are not available. To address this, we propose an Unsupervised Deep Video Denoiser (UDVD¹), a CNN architecture designed to be trained exclusively with noisy data. The performance of UDVD is comparable to the supervised state-of-the-art, even when trained only on a single short noisy video. We demonstrate the promise of our approach in real-world imaging applications by denoising raw video, fluorescence-microscopy and electron-microscopy data. In contrast to many current approaches to video denoising, UDVD does not require explicit motion compensation. This is advantageous because motion compensation is computationally expensive, and can be unreliable when the input data are noisy. A gradient-based analysis reveals that UDVD automatically adapts to local motion in the input noisy videos. Thus, the network learns to perform implicit motion compensation, even though it is only trained for denoising.

1. Introduction

Video denoising is a fundamental problem in image processing, as well as an important preprocessing step for computer vision tasks. Convolutional neural networks (CNNs) [21] provide current state-of-the-art solutions for this problem [34, 35, 41, 43, 11, 9, 8, 6]. These networks are typically trained using a database of clean videos, which are corrupted with simulated noise. However, in applications such as microscopy, noiseless ground truth videos are often not available. To address this issue, we propose a method to train a video denoising CNN without access to super-

vised data, which we call Unsupervised Deep Video Denoising (UDVD). UDVD is inspired by the “blind-spot” technique, recently introduced for unsupervised still image denoising [22, 17, 2, 19], in which a CNN is trained to estimate each *noisy* pixel from the surrounding spatial neighborhood *without including the pixel itself*. Here, we propose a blind-spot architecture that processes the surrounding spatio-temporal neighborhood to denoise videos.

We show that UDVD is competitive with the current supervised state-of-the-art on standard benchmarks, despite not having access to ground-truth clean videos during training (see Figure 1). Moreover, when combined with aggressive data augmentation and early stopping, it can produce high-quality denoising even when trained exclusively on a single *brief* noisy video sequence (as few as 30 frames), outperforming unsupervised video denoising techniques (e.g. F2F[11] and MF2F [9]) which are pre-trained with supervision. Finally, methods based on pre-training are not suitable for imaging applications where clean data is unavailable. In contrast, we demonstrate that UDVD can effectively denoise three different real-world datasets: raw videos from surveillance cameras, fluorescence-microscopy videos of cells, and electron-microscopy videos of catalytic nanoparticles.

The state-of-the-art performance of UDVD is unexpected. Nearly all existing approaches to video denoising [24, 1, 3, 25], including those based on deep CNNs [34, 41, 11, 13, 42], use estimates of optical flow to adaptively compensate for the motion of objects in the video. Conventional wisdom suggest that ignoring such motion should lead to denoising results in which moving content is blurred. Contrary to this intuition, UDVD and some recent state-of-the-art supervised methods for video denoising [35, 8, 6] yield excellent empirical performance without explicit estimation of optical flow. *How can is this achieved?* We use a gradient-based analysis to show that both UDVD and supervised CNNs perform spatio-temporal *adaptive* filter-

*equal contribution.

¹See <https://sreyas-mohan.github.io/udvd/> for code and more results.

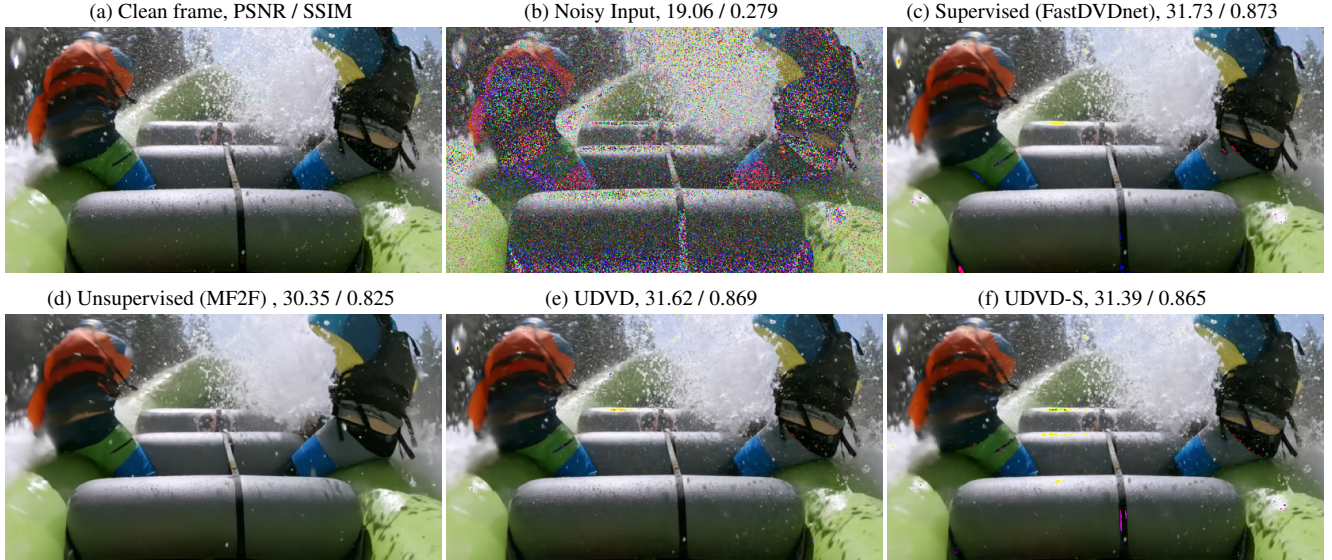


Figure 1. **Unsupervised denoising matches the performance of supervised denoising.** Frame from a video in the Set8 dataset denoised using different approaches. (a) Clean frame. (b) Frame corrupted with Gaussian noise of standard deviation 30 (relative to intensity range [0-255]). (c) FastDVDnet [35], a supervised method trained on the DAVIS dataset. (d) MF2F [9], an unsupervised method which fine-tunes a pre-trained FastDVDnet on the noisy video (e) Our proposed unsupervised method (UDVD), which uses five frames to denoise each frame, trained on the DAVIS dataset. (f) UDVD trained only on the noisy video itself. Performance is quantified using PSNR / SSIM [38], respectively. The corresponding videos, as well as additional examples, are included in Section C of the supplementary material.

ing, which is aligned with underlying motion. Thus, these CNNs are *automatically performing implicit motion compensation*. To quantify this, we demonstrate that it is possible to estimate optical flow accurately from the network gradients, even though the network architectures are not designed to account for optical flow, and the models receive no optical-flow information during training.

Our Contributions:

- A novel blind-spot architecture/objective for unsupervised video denoising, which achieves performance competitive with state-of-the-art supervised methods.
- A training paradigm using aggressive data augmentation (time and space reversal) and early stopping to achieve state-of-the-art performance from training *on a single brief noisy video*.
- A demonstration of our method’s effectiveness in denoising real-world electron and fluorescence microscopy data, as well as raw videos. Unlike most existing methods for unsupervised video denoising, our proposed method does not require pre-training, which is key in real-world imaging applications.
- An analysis of the denoising mechanism learned by UDVD, demonstrating that it performs implicit motion compensation even though it is only trained for denoising. We apply the analysis to supervised networks, showing that the same conclusion holds.

2. Background and Related Work

Traditional and CNN-based video denoising. Traditional techniques for single image denoising include nonlinear filtering [36, 26], sparse prior methods [12, 10, 33, 4, 30, 7], and nonlocal means [20]; many of which have been extended to videos [24, 1, 25, 3]. In order to exploit the spatio-temporal structure of the video, these methods typically employ motion compensation based on estimates of optical flow.

In the last five years, data-driven methods based on deep CNNs [21] have outperformed all other techniques in image [45, 14, 5] and video denoising [34, 41, 35, 43]. The CNNs are trained to minimize the mean squared error between the network output and ground truth using large databases of natural images/videos. Many deep-learning techniques also perform explicit motion compensation. DVDnet [34] applies an image-denoising CNN to each input frame, estimates the optical flow from the denoised frames using DeepFlow [39] (a CNN pre-trained for this purpose), warps the frames using the flow estimate to align their content, and finally processes the registered frames with a CNN. Ref. [41] applies a similar pipeline, but jointly trains an optical-flow module with the denoising CNN.

Video denoising without motion compensation. Three recent methods perform video denoising without explicit motion estimation. VNLnet [8] uses a non-local search algorithm to find self-similar patches in the input video, and then uses a CNN to process the patches. ViDeNN [6] consists of

a first stage that denoises each frame using a CNN, and a second stage that exploits temporal structure by using the frames, $(t - 1)$, t and $t + 1$ to produce the denoised t th frame. FastDVDnet [35] uses UNet [32] blocks, trained end to end, to denoise each frame using five contiguous frames. These methods achieve state-of-the-art performance without any explicit motion compensation, similar to our proposed UDVD. In this work we show that such CNNs actually performs *implicit* motion estimation, which can be revealed through a gradient-based analysis.

Unsupervised denoising. Noise2Noise (N2N) is an unsupervised image-denoising technique where a CNN is trained on pairs of noisy images corresponding to the same clean image [22]. Frame2Frame (F2F) [11] exploits this approach to fine-tune a pretrained image-denoising CNN with noisy data. The idea is to register contiguous frames using the optical flow (obtained from TV-L1 [44]), and treat them as noisy realizations of the same clean image. This scheme is extended to have a trainable flow estimation module in [42], additional optical-flow consistency in [13] and to use multiple noisy frames as input in Multi-Frame2Frame (MF2F) [9].

Using the N2N framework to perform unsupervised video denoising requires warping adjoining frames, which in turn requires explicit motion compensation, and accurate occlusion estimation. In addition, the assumption that contiguous frames can be registered may not hold, particularly if the motion speeds in the video are large relative to the frame rate or local intensity changes are not due to translation. In order to bypass these issues, we develop a blind-spot network that trains denoising CNNs by fitting the noisy data directly. The CNN is trained to estimate each noisy pixel value using the surrounding spatio-temporal neighborhood, but without taking into account the noisy pixel itself in order to avoid the trivial identity solution. This “blind spot” can be enforced through architecture design [19], or by masking [2, 17]. For still images, several variations of this approach have been shown to provide effective denoising for natural images and noisy images from fluorescence microscopy [18, 31, 15].

3. Unsupervised Deep Video Denoising

In this section we describe our proposed architecture (see Figure 2 for a detailed diagram).

Multi-frame blind-spot architecture. Our CNN maps five contiguous noisy frames to a denoised estimate of the middle frame. Building on the “blind spot” idea proposed in [19] for single-image denoising, we design the architecture so that each output pixel is estimated from a spatio-temporal neighbourhood that does not include the pixel itself. We rotate the input frames by multiples of 90° and process them through four separate branches containing asymmetric convolutional filters that are *vertically causal*. As a

result, the branches produce outputs that only depend on the pixels above (0° rotation), to the left (90°), below (180°) or to the right (270°) of the output pixel. These partial outputs are then *derotated* and combined using a three-layered cascade of 1×1 convolutions and nonlinearities to produce the final output. The resulting field of view does not include the pixel being denoised, as depicted at the bottom of Figure 2.

UDVD processes the video in two stages as shown in Figure 2, similar to previously proposed networks for supervised video denoising [34, 6, 35]. A first stage, consisting of three UNets [32] (D1 in the diagram) with shared parameters, maps each group of three contiguous frames (i.e. $(t - 2, t - 1, t)$, $(t - 1, t, t + 1)$ and $(t, t + 1, t + 2)$) to a separate feature map. These features are then mapped to a single output using another UNet (D2). See Suppl. A for a detailed description of the architecture.

Bias-free architecture. Inspired by [27], we remove all additive terms from the convolutional layers in UDVD. This provides automatic generalization to varying noise levels not encountered during training, and facilitates our proposed analysis to interpret the denoising mechanisms learned by the network (see Section 5 and 6).

Using the missing pixel. The denoised value generated by the proposed architecture at each pixel is computed without using the noisy observation at that location. This avoids overfitting – i.e. learning the trivial identity map that minimizes the mean-squared error cost function – but ignores important information provided by the noisy pixel. In the special case of Gaussian additive noise, we can use this information via a precision-weighted average between the network output and the noisy pixel value. Following [19, 18], the weights in the average are derived by assuming a Gaussian distribution for the error in the blind-spot estimates of the color pixel values. Specifically, we model the distribution of the three color channels of a pixel $x \in \mathcal{R}^3$ given the noisy neighbourhood Ω_y as $p(x|\Omega_y) = \mathcal{N}(\mu_x, \Sigma_x)$, where $\mu_x \in \mathcal{R}^3$ and $\Sigma_x \in \mathcal{R}^3$ represent the mean vector and covariance matrix. Let $y = x + \eta$, $\eta \sim \mathcal{N}(0, \sigma^2 I_3)$ be the observed noisy pixel. We integrate the information in the noisy pixel with the UDVD output by computing the mean of the posterior $p(x|y, \Omega_y)$, given by

$$E[x|y] = (\Sigma_x^{-1} + \sigma^{-2}I)^{-1}(\Sigma_x^{-1}\mu_x + \sigma^{-2}y). \quad (1)$$

See Suppl. A for more details. The CNN architecture is trained to estimate the mean and covariance of this distribution at each pixel by maximizing the log likelihood of the noisy data:

$$\begin{aligned} \mathcal{L}(\mu_x, \Sigma_x) = & \frac{1}{2}[(y - \mu_x)^T(\Sigma_x + \sigma^2 I)^{-1}(y - \mu_x)] \\ & + \frac{1}{2} \log|\Sigma_x + \sigma^2 I|. \end{aligned} \quad (2)$$

When the noise process is unknown, we simply minimize

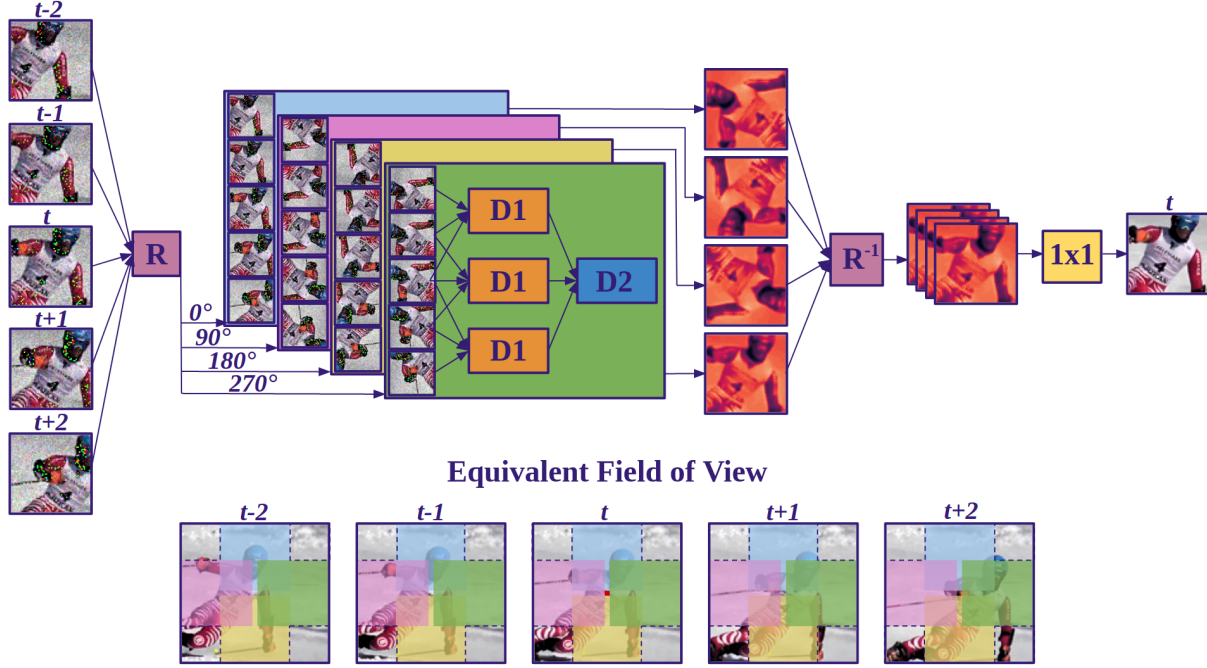


Figure 2. **Unsupervised Deep Video Denoising (UDVD) Network Architecture.** The network takes 5 consecutive noisy frames as input and produces a denoised central frame as output. We rotate the input frames by multiples of 90° and process them in four separate branches with shared parameters, each containing asymmetric convolutional filters that are *vertically causal*. As a result, the branches produce outputs that only depend on the pixels above (0° rotation, blue region), to the left (90° , pink region), below (180° , yellow region) or to the right (270° , green region) of the output pixel. Each branch consists of a cascade of 2 Unet-style blocks (D1 and D2) to combine information over frames. These outputs are then *derotated* and linearly combined (using a 1×1 convolutions) followed by a ReLU nonlinearity to produce the final output. The resulting “field of view” is depicted at the bottom with each color representing the contribution of the corresponding branch.

the MSE between the denoised output and noisy video, and ignore the center pixel (see Suppl. A for more details).

Data augmentation and early stopping. In supervised denoising with simulated noise, training can rely on the generation of a virtually unlimited set of fresh noise realizations, which prevents overfitting. In the unsupervised setting, this is not possible, which makes it more challenging to train models that can denoise short video sequences. To address this, we (a) leverage data augmentation strategies: spatial flipping and time reversal, and (b) perform early stopping by monitoring the mean squared error between the network output and noisy frames on a held-out set of frames. These strategies make it possible to train UDVD with short video sequences (as few as 30 frames), while achieving denoising performance that is on par with or superior to both unsupervised and supervised networks trained on much larger datasets (see Figure 1, Table 2 and Suppl. D).

4. Datasets

We demonstrate the broad applicability of our approach by validating it on domains with different signal and noise

structure: natural videos, raw videos, fluorescence microscopy, and electron microscopy.

Natural videos. We perform controlled experiments on natural videos by adding iid Gaussian noise to the DAVIS dataset [29]. The training/validation/test split is 60/30/30 videos, respectively. We use three additional datasets for testing - Set8 [35] composed of 4 videos from the Derfs Test Media collection and 4 videos captured with a GoPro camera, Derfs [9] with 7 videos, and the first 10 videos from Vid3oC [16] dataset (See Suppl. D for details).

Raw videos. We evaluate UDVD on a dataset of raw videos i.e with frame color channels interleaved according to the sensor mosaic containing real noise introduced in [43]. The dataset contains 11 unique videos, each containing 7 frames, captured at five different ISO levels using a surveillance camera. Each video has 10 different noise realizations per frame, which are averaged to obtain an estimated clean version of the video.

Fluorescence microscopy. We apply our approach to fluorescence-microscopy recordings of live cells in [37]. We use two videos: Fluo-C2DL-MSC (CTC-MSC) depicting mesenchymal stem cells, and Fluo-N2DH-GOWT1

test set	σ	Traditional		Supervised CNN			Unsupervised CNN (UDVD)		
		VNLB	VBM4D	VNLnet	DVDnet	FastDVDnet	1 frame	3 frames	5 frames
DAVIS	30	33.73	31.65	-	34.08	34.06	32.80	33.48	33.92
	40	32.32	30.05	32.32	32.86	32.80	31.48	32.20	32.68
	50	31.13	28.80	31.43	31.85	31.83	30.47	31.20	31.70
Set8	30	31.74	30.00	-	31.79	31.60	30.91	31.62	32.01
	40	30.39	28.48	30.55	30.55	30.37	29.63	30.42	30.82
	50	29.24	27.33	29.47	29.56	29.42	28.65	29.47	29.89

Table 1. **Denoising results on natural video datasets.** All networks are trained on the DAVIS train set. Performance values are PSNR of each trained network averaged over held-out test data. UDVD, operating on 5 frames, outperforms the supervised methods on Set8 and is competitive on the DAVIS test set. Unsupervised denoisers with more temporal frames show a consistent improvement in denoising performance. DVDnet and FastDVDnet are trained using varying noise levels ($\sigma \in [0, 55]$) and VNLnet is trained and evaluated on each specified noise level. All UDVD networks are trained *only* at $\sigma = 30$, showing that they generalize well on unseen noise levels. See Sections C and F in the supplementary material for additional results. The PSNR values for all methods except UDVD are taken from [35].

	$\sigma = 30$				$\sigma = 90$			
	DAVIS	Set8	Derfs	Vid3oC	DAVIS	Set8	Derfs	Vid3oC
UDVD-S	33.68 / 78.16	32.90 / 81.85	33.95 / 81.91	34.65 / 84.60	29.05 / 53.53	28.07 / 55.35	29.42 / 59.25	29.94 / 63.79
UDVD*	33.78 / 79.88	31.90 / 82.53	32.58 / 81.44	34.24 / 83.96	28.87 / 51.22	27.25 / 51.84	28.26 / 52.44	29.23 / 60.08
FastDVDnet*	33.91 / 76.99	31.81 / 80.21	32.45 / 81.64	35.05 / 84.44	28.01 / 47.53	26.54 / 50.16	27.36 / 52.87	28.42 / 55.99
MF2F	33.91 / 80.01	31.84 / 80.55	32.87 / 82.22	35.18 / 85.71	28.81 / 51.24	27.25 / 52.78	28.29 / 55.06	29.67 / 61.28

Table 2. **Results for UDVD trained on individual noisy videos.** The top row shows PSNR/VMAF[23] values (averaged over the entire dataset) for UDVD trained on each individual video sequence with early stopping (labelled UDVD-S) using the last 5 frames of a video as a held-out set. We augmented the dataset with spatial flipping and time reversal (see Suppl. D for an ablation study). With the augmentations and early stopping, UDVD-S is comparable to (and often outperforms) UDVD or FastDVDnet trained on the full DAVIS dataset (indicated by *) and MF2F, which fine-tunes a pre-trained CNN on each individual video. See Suppl. D for results on individual video sequences.

(CTC-N2DH) depicting GOWT1 cells. This dataset illustrates the challenges of applying supervised approaches to real data: there is no ground-truth clean data.

Electron microscopy. We also apply our methodology to a transmission electron microscopy dataset from [28]. The data consist of a 40-frame video depicting a platinum nanoparticle supported on a cerium oxide base. The average image intensity is 0.45 electrons/pixel, which results in an extremely low signal-to-noise ratio. As with the fluorescence-microscopy data, no ground-truth clean images are available.

5. Experiments and Results

Comparison with other approaches on natural videos. We train UDVD on the DAVIS training set (see Suppl. A for the training procedure). Following [45, 27, 35, 34, 19, 17, 2], we add iid Gaussian noise with standard deviation $\sigma = 30$ on the clean videos during training. UDVD is evaluated on the DAVIS test set and on Set8 by comparing to the clean ground-truth videos via PSNR. We compare UDVD with several popular methods: Bayesian processing of spatio-temporal patches (VNLB [20]), an extension of the popular image-denoising algorithm BM3D (VBM4D [25]) and supervised CNNs (VNLnet [8], DVD-

net [34], FastDVDnet [35]). As shown in Table 1, UDVD achieves comparable performance to the supervised state-of-the-art on the DAVIS test set and slightly outperforms these methods on an independent test set (Set8) at multiple noise levels. It also outperforms traditional unsupervised techniques such as VNLB and VBM4D (see Figure 1 and Suppl. C for visual examples).

Unsupervised denoising from limited data. In order to validate our approach on a more challenging setting that is closer to the practical applications of unsupervised denoising, we trained and tested UDVD on individual videos from our test sets. As shown in Table 3 and 4 in Suppl. D, when combined with data augmentation and early stopping (using the last 5 frames of each video as a held-out validation set), this version of UDVD (called UDVD-S) achieves comparable results, or often outperforms supervised FastDVDnet and unsupervised UDVD trained on a large dataset (DAVIS) (see Table 2 for results on 4 different datasets).

To the best of our knowledge, all the existing unsupervised video denoising techniques are based on the F2F [11] framework, where a backbone CNN pre-trained with supervision is fine-tuned on the video to be denoised. We compared UDVD-S against the most recent such method – MF2F [9] which fine-tunes a FastDVDnet [35] trained

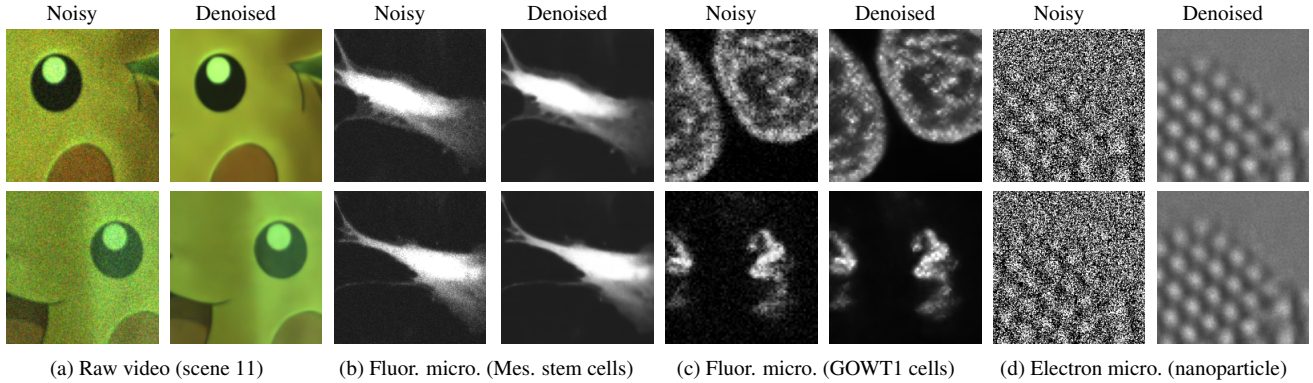


Figure 3. **Denoising real-world data.** Results from applying UDVD to the raw video, fluorescence-microscopy and electron-microscopy datasets described in Section 4. Qualitatively, UDVD succeeds in removing noise while preserving the underlying signal structure, even for the highly noisy electron-microscopy data. Raw videos are converted to RGB for visualization. See Suppl. D and F for denoised videos.

with supervision on natural videos using an objective involving optical flow computed on consecutive noisy frames (see Section 2). Without any pre-training, UDVD-S outperforms MF2F in almost all videos in Table 3 and 4 in Suppl. D, and datasets in Table 2 (See Table 5 in Suppl. D.3 for measure of confidence). Note that (a) we trained MF2F using all the 5 training schemes provided in the paper and reported the best results in Table 2, and (b) the metric we used to measure performance in Table 2 is the average PSNR of all denoised frames, unlike in Ref. [9] where the first 10 frames of each video were excluded (see Suppl. D.3 for more details and results).

Use of temporal information. UDVD estimates each frame from k surrounding contiguous frames. To validate the effect of using more temporal information, we tested $k \in \{1, 3, 5\}$. As shown in Table 1, performance improves substantially and monotonically with k (see Suppl. B for more noise levels). This is in agreement with the literature on supervised learning [35]. The performance gains arising from a longer temporal context are more substantial at higher noise levels (see Table 1). This is consistent with our analysis in Section 6 which shows that, at low noise levels, UDVD($k = 5$) tends to ignore the distant frames, but relies on them more at higher noise levels (see Figure 4 & Suppl. G).

Generalization across noise levels. UDVD generalizes strongly across noise levels not encountered during training. The results in Table 1 are obtained with a network trained only at a fixed noise level of $\sigma = 30$. This generalization ability is consistent with bias-free networks for image denoising [27]. See Suppl. F for more discussion and results.

Raw videos with real noise. We train UDVD on the first 9 realizations of the 5 videos from the test set of the raw video dataset (see Section 4), holding out the last realization for early stopping. We compare our performance with RViDeNet [43] which is pre-trained on a simulated dataset

CNN \ ISO	1600	3200	6400	12800	25600	mean
UDVD	48.04	46.24	44.70	42.19	42.11	44.69
RViDeNet [43]	47.74	45.91	43.85	41.20	41.17	43.97

Table 3. **Raw video denoising.** PSNR values evaluated on the test set of the raw video dataset (Section 4) when denoised with (a) UDVD trained only the noisy test videos and (b) RViDeNet trained with supervision on a large dataset. The columns correspond to different ISO levels, with larger levels resulting in noisier data.

and then fine-tuned with supervision on 6 training videos from the raw video dataset. UDVD outperforms RViDeNet at all noise levels (see Table 3 and Fig 3). Note that UDVD was directly trained on the mosaiced raw videos. Existing unsupervised video denoising methods, like MF2F, cannot be applied directly on this dataset as their pre-trained backbone expects an input in the RGB domain (more details in Suppl. E).

Real-world microscopy data. We train UDVD on the fluorescence-microscopy data described in Section 4 following the same procedure as for the natural videos, including data augmentation. For the electron-microscopy data, we trained on the first 35 frames of the video, and used the remaining 5 as a validation set to perform early stopping based on mean-squared error. UDVD is able to effectively denoise the fluorescence-microscopy and the electron-microscopy datasets described in Section 4. This can be appreciated qualitatively in Figure 3 and Suppl. E.

6. Automatic Motion Compensation

Most previous approaches for video denoising rely on explicit motion compensation [24, 1, 3, 25]. This requires estimating the optical flow, which is the local translational motion of features in the image arising from the motion of objects and surfaces in a visual scene relative to the cam-

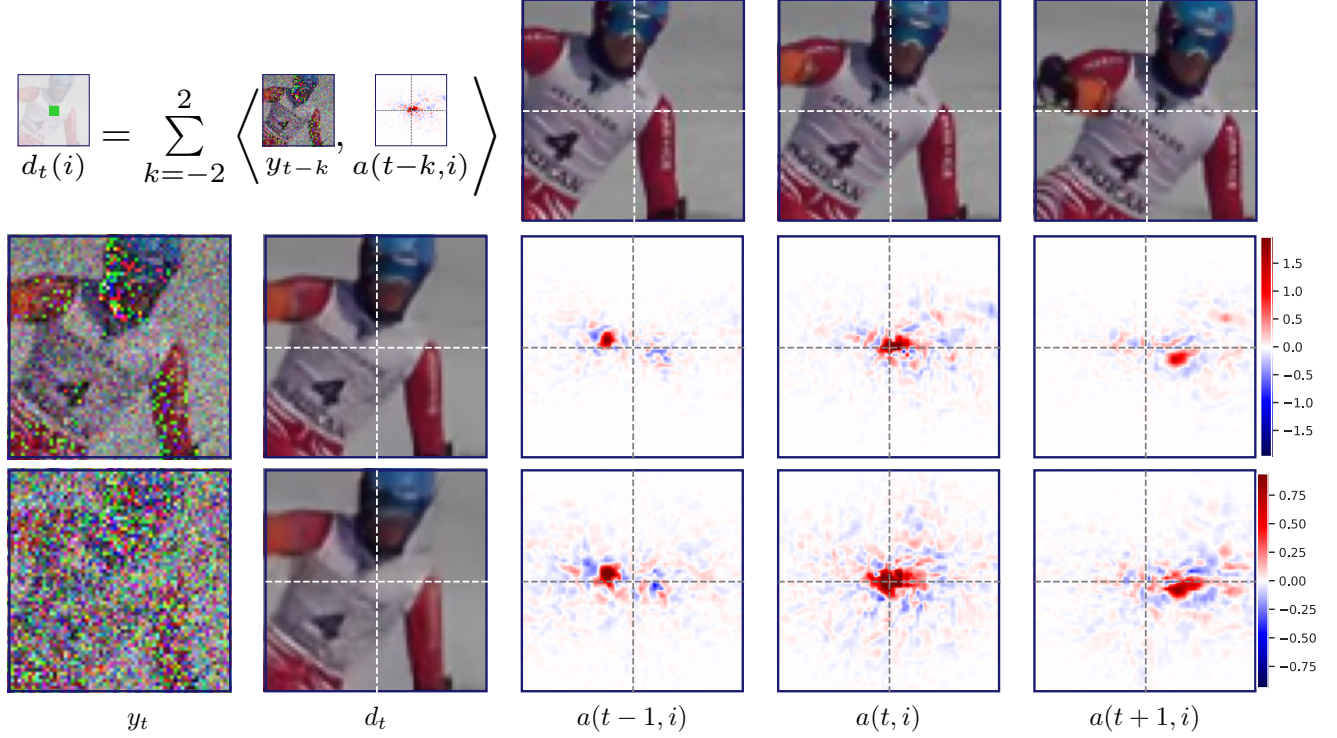


Figure 4. **Video denoising as spatiotemporal adaptive filtering.** Visualization of the equivalent linear weights ($a(k, i)$, Eq. 4) used to compute two example denoised pixels using UDVD. The left two columns show noisy frames y_t at two noise levels, and the corresponding denoised frames, d_t . Three successive clean frames $\{x_{t-1}, x_t, x_{t+1}\}$ are shown in top row, for reference. Corresponding weights $a(k, i)$ for pixel i (intersection of the dashed white lines) in these three frames, are shown in the last three columns. The weights are seen to adapt to underlying video content, with their mode shifting to track the motion of the skier. As the noise level σ increases (bottom row), their spatial extent grows, averaging out more of the noise while respecting object boundaries. For each denoised pixel, the sum of weights (over all pixel locations and frames) is approximately one, and thus can be interpreted as computing a local average (but note that some weights are negative, depicted in blue).

era. Several CNN-based denoisers build motion estimation into the architecture [34, 41]. In particular, motion compensation is critical to the F2F and MF2F frameworks for unsupervised denoising, which use motion compensation to register contiguous images [11, 13, 9]. In contrast, recent supervised video denoising networks like FastDVDnet [35] and ViDeNN [6], as well as our unsupervised UDVD, do not perform any explicit motion compensation. Despite this, they achieve state-of-the-art results. The empirical performance of these approaches suggests that the networks must somehow be exploiting temporal information successfully. Here, we study this phenomenon through an analysis of the denoising mapping, which reveals that these networks perform an implicit form of motion compensation.

Gradient-based analysis. We use the approach of [27] to analyze CNNs trained for image denoising. Let $y \in \mathbb{R}^{nT}$ be a flattened video sequence containing T noisy frames with n pixels each, processed by a CNN. We define the denoising function $f_i : \mathbb{R}^{nT} \rightarrow \mathbb{R}$ as the map between the noisy video and the denoised value $d_i := f_i(y)$ of the CNN output at the i th pixel. A first-order Taylor decomposition of

the denoising function may be written as:

$$d_i := f_i(y) = \langle \nabla f_i(y), y \rangle + b, \quad (3)$$

where $\nabla f_i(y) \in \mathbb{R}^{nT}$ denotes the gradient of f_i at y . The constant $b := f_i(y) - \langle \nabla f_i(y), y \rangle$ is the net bias of the network, a combined function of all additive constants in the convolutional and batch-normalization layers of the CNN.

Our proposed architecture is bias-free (i.e., all additive constants are removed from the architecture, as proposed in [27]), and thus $b = 0$. As a result, the denoised value at the i th pixel may be written as:

$$d(i) = \langle \nabla f_i(y), y \rangle = \sum_{k=1}^T \langle a(k, i), y_k \rangle, \quad (4)$$

where y_k denotes each of the T flattened frames that compose the noisy video, and the weights $a(k, i)$ correspond to the gradient of f_i with respect to y . Each vector $a(k, i)$ can be interpreted as an *equivalent filter* that produces an estimate of the denoised video at pixel i via a weighted average of the noisy observations over space and time.

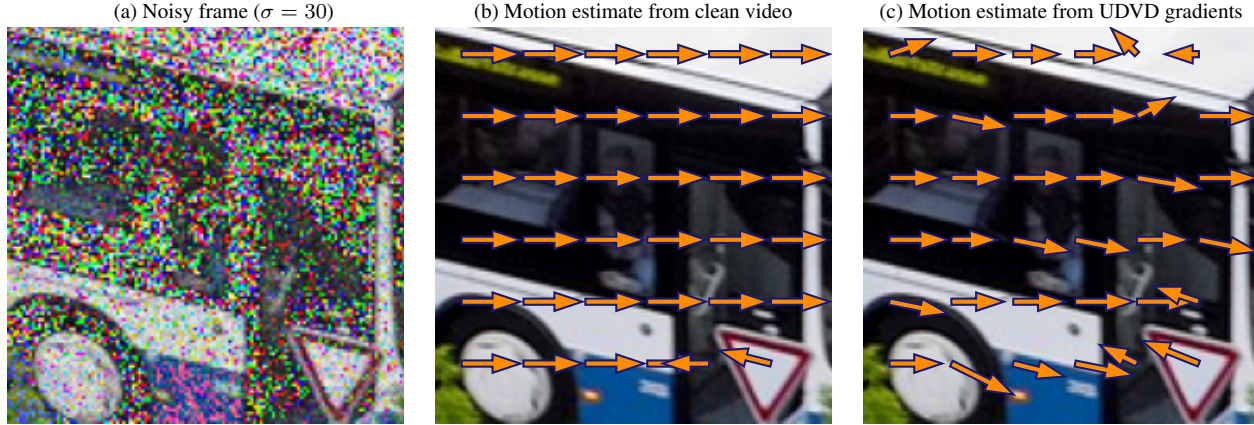


Figure 5. **CNNs trained for denoising automatically learn to perform motion estimation.** (a) Noisy frame from a video in the DAVIS dataset. (b) Optical flow direction at multiple locations of the image obtained using a state-of-the-art algorithm applied to the *clean video*. (c) Optical flow direction estimated from the shift of the adaptive filter obtained by differentiating the network, which is trained exclusively with noisy videos and no optical flow information. Optical flow estimates are well-matched to those in (b), but deviate according to the aperture problem at oriented features (see black vertical edge of bus door), and in homogeneous regions (see bus roof, top right).

Interpreting equivalent filters. Visualizing these equivalent filters reveals that UDVD learns to denoise by performing averaging over an adaptive spatiotemporal neighborhood of each pixel. As illustrated in Figure 4 (and Suppl. G), when the noise level increases, the averaging is carried out over larger regions. This intuitive behavior is also seen in classical linear Wiener filters [40], where the filters are larger for higher levels of noise. The crucial difference is that in the case of CNNs, the equivalent filters are *adapted* to the local video content: they respect object boundaries in space and time, taking into account their motion. This is apparent in Figure 4: equivalent filters in adjoining frames are automatically shifted spatially to compensate for the movement of the skier (additional examples in Suppl. G). We find that this implicit motion compensation is not unique to UDVD: CNNs trained in a supervised fashion have the same property (see also Suppl. G).

Optical-flow estimation. In order to validate our observation that CNNs exclusively trained for denoising implicitly detect and exploit video motion, we use the equivalent filters of the networks to estimate the optical flow. To estimate the optical flow from the t^{th} frame to the $(t+1)^{\text{th}}$ frame at the i^{th} pixel, we compute the difference between the position of the centroid of the equivalent filter corresponding to the pixel at times t , $a(t, i)$, and time $t+1$, $a(t+1, i)$. To increase the stability of the estimated flow, we compute the filter centroid through a robust weighted average that only includes entries with relatively large values (within 20% of maximum value in the filter).

The optical-flow estimates obtained from the gradients of the trained UDVD network are surprisingly precise, even at very high noise levels. Figure 5, and additional figures in Suppl. G, show that the results are similar to those obtained

by applying an algorithm for optical-flow estimation (DeepFlow [39]) on the corresponding *clean video*. This demonstrates that the CNNs are able to implicitly estimate motion from data, despite the fact that they were not trained on that problem, and *even in the presence of substantial noise corruption*, a setting that is quite challenging for optical-flow estimation techniques. We also observe that the optical-flow estimates obtained from UDVD gradients tend to be less accurate for pixels near strongly oriented features where local motion is only partially constrained (known as the *aperture problem*) or in homogeneous regions, where the local motion is unconstrained (the *blank wall problem*).

7. Conclusion

In this work we propose a method for unsupervised deep video denoising that achieves comparable performance to state-of-the-art supervised approaches. Combined with data-augmentation techniques and early stopping, the method achieves effective denoising even when trained exclusively on short individual noisy sequences, which enables its application to real-world noisy data. In addition, we perform a gradient-based analysis of denoising CNNs, which reveals that they learn to perform implicit adaptive motion compensation. This suggests several interesting research directions. For example, denoising may be a useful pretraining task for optical-flow estimation or other computer-vision tasks requiring motion estimation.

Acknowledgements. This work was supported by HHMI, NSF NRT HDR Award 1922658, CBET 1604971, OAC-1940263 and OAC-1940097. We thank the HPC staff at NYU, ASU and RBCDSAI, IIT Madras for their support.

References

- [1] Pablo Arias and Jean-Michel Morel. Video denoising via empirical Bayesian estimation of space-time patches. *Journal of Mathematical Imaging and Vision*, 60(1):70–93, 2018. 1, 2, 6
- [2] Joshua Batson and Loic Royer. Noise2self: Blind denoising by self-supervision. In *Proceedings of the 36th International Conference on Machine Learning*, pages 524–533, 2019. 1, 3, 5
- [3] Antoni Buades, Jose-Luis Lisani, and Marko Miladinović. Patch-based video denoising with optical flow estimation. *IEEE Transactions on Image Processing*, 25(6):2573–2586, 2016. 1, 2, 6
- [4] S Grace Chang, Bin Yu, and Martin Vetterli. Adaptive wavelet thresholding for image denoising and compression. *IEEE Trans. Image Processing*, 9(9):1532–1546, 2000. 2
- [5] Yunjin Chen and Thomas Pock. Trainable nonlinear reaction diffusion: A flexible framework for fast and effective image restoration. *IEEE transactions on pattern analysis and machine intelligence*, 39(6):1256–1272, 2016. 2
- [6] Michele Claus and Jan van Gemert. Videnn: Deep blind video denoising. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshop*, pages 1843–1852, 2019. 1, 2, 3, 7
- [7] Kostadin Dabov, Alessandro Foi, Vladimir Katkovnik, and Karen Egiazarian. Image denoising by sparse 3-d transform-domain collaborative filtering. *IEEE Transactions on Image Processing*, pages 2080–2095, 2017. 2
- [8] Axel Davy, Thibaud Ehret, Jean-Michel Morel, Pablo Arias, and Gabriele Facciolo. A non-local CNN for video denoising. In *2019 IEEE International Conference on Image Processing (ICIP)*, pages 2409–2413. IEEE, 2019. 1, 2, 5
- [9] Valery Dewil, Jeremy Anger, Axel Davy, Thibaud Ehret, Gabriele Facciolo, and Pablo Arias. Self-supervised training for blind multi-frame video denoising. In *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision*, pages 2724–2734, 2021. 1, 2, 3, 4, 5, 6, 7
- [10] D Donoho. Denoising by soft-thresholding. *IEEE Trans Info Theory*, 43:613–627, 1995. 2
- [11] Thibaud Ehret, Axel Davy, Jean-Michel Morel, Gabriele Facciolo, and Pablo Arias. Model-blind video denoising via frame-to-frame training. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 11361–11370, 2019. 1, 3, 5, 7
- [12] Michael Elad and Michal Aharon. Image denoising via sparse and redundant representations over learned dictionaries. *IEEE Transactions on Image processing*, 15(12):3736–3745, 2006. 2
- [13] Bichuan Guo, Jiangtao Wen, Zhen Xia, Shan Liu, and Yuxing Han. Learning model-blind temporal denoisers without ground truths. *arXiv preprint arXiv:2007.03241*, 2020. 1, 3, 7
- [14] Shi Guo, Zifei Yan, Kai Zhang, Wangmeng Zuo, and Lei Zhang. Toward convolutional blind denoising of real photographs. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 1712–1722, 2019. 2
- [15] Wesley Khademi, Sonia Rao, Clare Minnerath, Guy Hagen, and Jonathan Ventura. Self-supervised Poisson-Gaussian denoising. In *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision*, pages 2131–2139, 2021. 3
- [16] Sohyeong Kim, Guanju Li, Dario Fuoli, Martin Danelljan, Zhiwu Huang, Shuhang Gu, and Radu Timofte. The vid3oc and intvid datasets for video super resolution and quality mapping. In *2019 IEEE/CVF International Conference on Computer Vision Workshop (ICCVW)*, pages 3609–3616, 2019. 4
- [17] Alexander Krull, Tim-Oliver Buchholz, and Florian Jug. Noise2void - learning denoising from single noisy images. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 2124–2132, 2019. 1, 3, 5
- [18] Alexander Krull, Tomas Vicar, and Florian Jug. Probabilistic noise2void: Unsupervised content-aware denoising. *arXiv preprint arXiv:1906.00651*, 2019. 3
- [19] Samuli Laine, Tero Karras, Jaakko Lehtinen, and Timo Aila. High-quality self-supervised deep image denoising. In *Advances in Neural Information Processing Systems 32*, pages 6970–6980, 2019. 1, 3, 5
- [20] Marc Lebrun, Antoni Buades, and Jean-Michel Morel. A nonlocal Bayesian image denoising algorithm. *SIAM Journal on Imaging Sciences*, 6(3):1665–1688, 2013. 2, 5
- [21] Yann LeCun, Yoshua Bengio, and Geoffrey Hinton. Deep learning. *nature*, 521(7553):436, 2015. 1, 2
- [22] Jaakko Lehtinen, Jacob Munkberg, Jon Hasselgren, Samuli Laine, Tero Karras, Miika Aittala, and Timo Aila. Noise2noise: Learning image restoration without clean data. In *Proceedings of the 35th International Conference on Machine Learning*, pages 2965–2974, 2018. 1, 3
- [23] Zhi Li, Anne Aaron, Ioannis Katsavounidis, Anush Moorthy, and Megha Manohara. Toward a practical perceptual video quality metric. *Netflix Technology Blog*, 2016. 5
- [24] Ce Liu and William T Freeman. A high-quality video denoising algorithm based on reliable motion estimation. In *European conference on computer vision*, pages 706–719. Springer, 2010. 1, 2, 6
- [25] Matteo Maggioni, Giacomo Boracchi, Alessandro Foi, and Karen Egiazarian. Video denoising, deblocking, and enhancement through separable 4-d nonlocal spatiotemporal transforms. *IEEE Transactions on image processing*, 21(9):3952–3966, 2012. 1, 2, 5, 6
- [26] Peyman Milanfar. A tour of modern image filtering: New insights and methods, both practical and theoretical. *IEEE signal processing magazine*, 30(1):106–128, 2012. 2
- [27] Sreyas Mohan, Zahra Kadhodaie, Eero P. Simoncelli, and Carlos Fernandez-Granda. Robust and interpretable blind image denoising via bias-free convolutional neural networks. In *Proceedings of the International Conference on Learning Representations*, 2020. 3, 5, 6, 7
- [28] Sreyas Mohan, Ramon Manzorro, Joshua L Vincent, Binh Tang, Dev Yashpal Sheth, Eero P Simoncelli, David S Matteson, Peter A Crozier, and Carlos Fernandez-Granda. Deep denoising for scientific discovery: A case study in electron microscopy. *arXiv preprint arXiv:2010.12970*, 2020. 5

- [29] Jordi Pont-Tuset, Federico Perazzi, Sergi Caelles, Pablo Arbeláez, Alex Sorkine-Hornung, and Luc Van Gool. The 2017 davis challenge on video object segmentation. *arXiv preprint arXiv:1704.00675*, 2017. 4
- [30] Javier Portilla, Vasily Strela, Martin J Wainwright, and Eero P Simoncelli. Image denoising using scale mixtures of Gaussians in the wavelet domain. *IEEE Trans. Image Processing*, 12(11), 2003. 2
- [31] Mangal Prakash, Manan Lalit, Pavel Tomancak, Alexander Krull, and Florian Jug. Fully unsupervised probabilistic noise2void. *arXiv preprint arXiv:1911.12291*, 2019. 3
- [32] Olaf Ronneberger, Philipp Fischer, and Thomas Brox. U-net: Convolutional networks for biomedical image segmentation. In *International Conference on Medical image computing and computer-assisted intervention*, pages 234–241. Springer, 2015. 3
- [33] E P Simoncelli and E H Adelson. Noise removal via Bayesian wavelet coring. In *Proc 3rd IEEE Int’l Conf on Image Proc*, volume I, pages 379–382, Lausanne, Sep 16-19 1996. IEEE Sig Proc Society. 2
- [34] Matias Tassano, Julie Delon, and Thomas Veit. Dvdnet: A fast network for deep video denoising. In *Proceedings of the IEEE International Conference on Image Processing*, pages 1805–1809, 2020. 1, 2, 3, 5, 7
- [35] Matias Tassano, Julie Delon, and Thomas Veit. Fastdvdnet: Towards real-time deep video denoising without flow estimation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 1351–1360, 2020. 1, 2, 3, 4, 5, 6, 7
- [36] Carlo Tomasi and Roberto Manduchi. Bilateral filtering for gray and color images. In *ICCV*, volume 98, 1998. 2
- [37] Vladimír Ulman et al. An objective comparison of cell-tracking algorithms. *Nature Methods*, 14:1141–1152, 2017. 4
- [38] Zhou Wang, Alan C Bovik, Hamid R Sheikh, Eero P Simoncelli, et al. Image quality assessment: from error visibility to structural similarity. *IEEE Trans. Image Processing*, 13(4):600–612, 2004. 2
- [39] Philippe Weinzaepfel, Jerome Revaud, Zaid Harchaoui, and Cordelia Schmid. Deepflow: Large displacement optical flow with deep matching. In *Proceedings of the IEEE international conference on computer vision*, pages 1385–1392, 2013. 2, 8
- [40] Norbert Wiener. *Extrapolation, interpolation, and smoothing of stationary time series: with engineering applications*. Technology Press, 1950. 8
- [41] Tianfan Xue, Baian Chen, Jiajun Wu, Donglai Wei, and William T Freeman. Video enhancement with task-oriented flow. *International Journal of Computer Vision (IJCV)*, 127(8):1106–1125, 2019. 1, 2, 7
- [42] Songhyun Yu, Bumjun Park, Junwoo Park, and Jechang Jeong. Joint learning of blind video denoising and optical flow estimation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) Workshops*, June 2020. 1, 3
- [43] H. Yue, C. Cao, L. Liao, R. Chu, and J. Yang. Supervised raw video denoising with a benchmark dataset on dynamic scenes. In *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 2298–2307, 2020. 1, 2, 4, 6
- [44] Christopher Zach, Thomas Pock, and Horst Bischof. A duality based approach for realtime tv-l 1 optical flow. In *Joint pattern recognition symposium*, pages 214–223. Springer, 2007. 3
- [45] Kai Zhang, Wangmeng Zuo, Yunjin Chen, Deyu Meng, and Lei Zhang. Beyond a Gaussian denoiser: Residual learning of deep CNN for image denoising. *IEEE Transactions on Image Processing*, pages 3142–3155, 2017. 2, 5