

Advancing Image Understanding in Poor Visibility Environments: A Collective Benchmark Study

Wenhan Yang*, Ye Yuan*, Wenqi Ren, Jiaying Liu, Walter J. Scheirer, Zhangyang Wang, Taiheng Zhang, Qiaoyong Zhong, Di Xie, Shiliang Pu, Yuqiang Zheng, Yanyun Qu, Yuhong Xie, Liang Chen, Zhonghao Li, and Chen Hong, Hao Jiang, Siyuan Yang, Yan Liu, Xiaochao Qu, Pengfei Wan, Shuai Zheng, Minhui Zhong, Taiyi Su, Lingzhi He, Yandong Guo, Yao Zhao, Zhenfeng Zhu, Jinxiu Liang, Jingwen Wang, Tianyi Chen, Yuhui Quan, Yong Xu, Bo Liu, Xin Liu, Qi Sun, Tingyu Lin, Xiaochuan Li, Feng Lu, Lin Gu, Shengdi Zhou, Cong Cao, Shifeng Zhang, Cheng Chi, Chubin Zhuang, Zhen Lei, Stan Z. Li, Shizheng Wang, Ruizhe Liu, Dong Yi, Zheming Zuo, Jianning Chi, Huan Wang, Kai Wang, Yixiu Liu, Xingyu Gao, Zhenyu Chen, Chang Guo, Yongzhou Li, Huicai Zhong, Jing Huang, Heng Guo, Jianfei Yang, Wenjuan Liao, Jiangang Yang, Liguozhou, Mingyue Feng, Likun Qin

Abstract—Existing enhancement methods are empirically expected to help the high-level end computer vision task; however, that is observed to not always be the case in practice. We focus on object or face detection in poor visibility enhancements caused by bad weathers (haze, rain) and low light conditions. To provide a more thorough examination and fair comparison, we introduce three benchmark sets collected in real-world hazy, rainy, and low-light conditions, respectively, with annotated objects/faces. We launched the UG²⁺ challenge Track 2 competition in IEEE CVPR 2019, aiming to evoke a comprehensive discussion and exploration about whether and how low-level vision techniques can benefit the high-level automatic visual recognition in various scenarios. To our best knowledge, this is the first and currently largest effort of its kind. Baseline results by cascading existing enhancement and detection models are reported, indicating the highly challenging nature of our new data as well as the large room for further technical innovations. Thanks to a large participation from the research community, we are able to analyze representative team solutions, striving to better identify the strengths and limitations of existing mindsets as well as the future directions.

Index Terms—Poor visibility environment, object detection, face detection, haze, rain, low-light conditions

*The first two authors Wenhan Yang and Ye Yuan contributed equally. Ye Yuan and Wenhan Yang helped prepare the dataset proposed for the UG2+ Challenges, and were the main responsible members for UG2+ Challenge 2019 (Track 2) platform setup and technical support. Wenqi Ren, Jiaying Liu, Walter J. Scheirer, and Zhangyang Wang were the main organizers of the challenge and helped prepare the dataset, raise sponsors, set up evaluation environment, and improve the technical submission. Other authors are the group members of winner teams in UG2+ challenge Track 2 contributing to the winning methods.

Wenhan Yang and Jiaying Liu are with the Wangxuan Institute of Computer Technology, Peking University, Beijing 100080, China (e-mail: yangwenhan@pku.edu.cn; liujiaying@pku.edu.cn).

Ye Yuan and Zhangyang Wang are with the Department of Computer Science and Engineering, Texas A&M University, TX 77843 USA (e-mail: atlaswang@tamu.edu; ye.yuan@tamu.edu).

Wenqi Ren is with State Key Laboratory of Information Security, Institute of Information Engineering, Chinese Academy of Sciences (e-mail: rwq.renwenqi@gmail.com).

W. Scheirer is with the Department of Computer Science and Engineering, University of Notre Dame, Notre Dame, IN, 46556 (e-mail: walter.scheirer@nd.edu).

Correspondence should be addressed to: Zhangyang Wang (atlaswang@tamu.edu), and Walter J. Scheirer (walter.scheirer@nd.edu)

I. INTRODUCTION

The arrival of the big data era brings us mass diverse applications and spawns a series of demands in both human and machine visions. On one hand, new applications in consumer electronics [1], such as TV broadcasting, movies, video-on-demand, *etc.*, expect continuous efforts to improve the human visual experience. On the other hand, many applications of smart cities and Internet of things (IoT), such as surveillance video, autonomous/assisted driving, unmanned aerial vehicles (UAVs), *etc.*, call for more effective and stable machine vision-based sensing and understanding [2]. As a result, it is a critical issue to explore the general framework that can benefit two kinds of tasks simultaneously and make them mutually beneficial.

However, most of the existing researches still aim to solve the problems in their own routes separately. In early researches [3–5], their models are not capable enough to consider beyond their own purposes (only for one of human vision or machine vision). Even in the recent decade, most of the enhancement methods, *e.g.*, dehazing [6–9], deraining [10–19], illumination enhancement [20–25], image/video compression [26–28], compression artifact removal [29, 30], and copy-move forgery detection [31], only target human vision and most of image/video understanding and analytics methods, *e.g.* classification [32], segmentation [33], action recognition [34], and human pose estimation [35], are considered to take the clean and high-quality images as the input of a system.

When the improved computation power and the new emerging data-driven approaches push for the progress of applying the existing state-of-the-art methods to industrial trials, lack of consideration on human and machine vision jointly leads to the observed fragility of real systems. Taking autonomous driving as an example: the industry players have been tackling the challenges posed by inclement weathers; However, heavy rain, haze or snow will still obscure the vision of on-board cameras and create confusing reflections and glare, leaving the state-of-the-art self-driving cars in the struggle [36]. Another illustrative example can be found in city surveillance: even the commercialized cameras adopted by governments appear

fragile in challenging weather conditions [37].

The largely jeopardized performance of visual sensing is caused by two aspects: *inconsistency between training and testing data*, *inappropriate guidance for network training*. First, most current vision systems are designed to perform in clear environments but the real-world scenes are unconstrained and might include dynamic degradation, *e.g.* moving platforms, bad weathers, and poor illumination, which are not covered in the training data for pretrained models. Second, the existing data-driven methods largely rely on task-driven tuning and extract task-related features. Therefore, they are sensitive to unseen contents and conditions. Once the model faces different degradation to the trained one or has a different target, *e.g.* switching from human vision to machine vision, the performance of the pretrained model might degrade much.

To face the real world in practical applications, a dependable vision system must reckon with the entire spectrum of complex unconstrained outdoor environments, instead of just working in limited scenes. However, at this point, existing academia and industrial solutions see significant gaps from addressing the above-mentioned pressing real-world challenges, and a systematic consideration and collective effort for identifying and resolving those bottlenecks that they commonly face have also been absent. Considering that, it is highly desirable to study to what extent, and in what sense, such challenging visual conditions can be coped with, for the goal of achieving robust visual sensing and understanding in the wild, which benefits security/safety, autonomous driving, robotics, and an even broader range of signal/image processing applications.

One primary challenge arises from the **Data** aspect. Those challenging visual conditions usually give rise to nonlinear and data-dependent degradations that will be much more complicated than the well-studied noise or motion blur. The state-of-the-art deep learning methods are typically hungry for training data. The usage of synthetic training data has been prevailing, but may inevitably lead to domain shifts [38]. Fortunately, those degradations often follow some parameterized physical models. That will naturally motivate a combination of model-based and data-driven approaches. In addition to training, the lack of real-world test sets (and consequently, the usage of potentially oversimplified synthetic sets) has limited the practical scope of the developed algorithms.

The other main challenge is found in the **Goal** side. Most restoration or enhancement methods cast the handling of those challenging conditions as a post-processing step of signal restoration or enhancement after sensing, and then feed the restored data for visual understanding. The performance of high-level visual understanding tasks will thus largely depend on the quality of restoration or enhancement. Yet it remains questionable whether restoration-based approaches would actually boost the visual understanding performance, as the restoration/enhancement step is not optimized towards the target task and may bring in misleading information and artifacts too. For example, a recent line of researches [8, 39–48] discuss on the intrinsic interplay relationship of low-level vision and high-level recognition/detection tasks, showing that their goals are not always aligned.

UG²⁺ Challenge Track 2 aims to evaluate and advance

the robustness of object detection algorithms in specific poor-visibility situations, including challenging weather and lighting conditions. We structure Challenge 2 into three sub-challenges. Each challenge features a different poor-visibility outdoor condition, and diverse training protocols (paired versus unpaired images, annotated versus unannotated, *etc.*). For each sub-challenge, we collect a new benchmark dataset captured in realistic poor-visibility environments with real image artifacts caused by rain, haze, insufficiency of light. The specific dataset details and evaluation protocols are illustrated in Section III. Comparing with previous works and challenges, our challenge and datasets include the following new features:

- **Covering Complex Degradation.** Our datasets capture images with synthetic and real haze (Challenge 2.1), under-exposure (Challenge 2.2), and rain streaks and rain-drops (Challenge 2.3), which provides precious resources to measure the statistics and properties of captured images in these scenes and design effective methods to recover the clean version of these images.
- **Supporting Un/Semi/Full-Supervised Learning.** In Challenge 2.1 and 2.2, our proposed datasets include paired and unpaired data to support both full-supervised or semi-supervised (optional for participants) training. In Challenge 2.3, there is no training data and testing samples provided. Therefore, it is a zero-shot problem and close to unsupervised learning. Therefore, our challenge supports all kinds of learning and provides important materials for future researches.
- **Highly Challenging.** In Challenge 2.1 and 2.2, the winners only achieve the results below 65 MAP. In Challenge 2.3, no participants achieve the results superior to the baseline results. The results show that, our datasets are challenging and there is still large room for further improvement.
- **Full Purpose.** Three sub-challenges support to evaluate the high-level tasks for machine vision. The Challenge 2.1 and 2.2 also include the paired data, which can support to evaluate for human vision.

The rest of the article is organized as follows. Section II briefly reviews previous works on poor visibility enhancement and visual recognition under adverse conditions as well as the related dataset efforts. Section III provides the detailed introduction of our datasets, challenges, evaluation protocols and baseline results. Section V illustrates the competition results and related analysis. Section VI summarizes the interesting observations, the reflected insights and briefly discusses the future directions. The concluding remarks are provided in Section VII.

II. RELATED WORK

A. Datasets

Most datasets used for image enhancement/processing mainly targets at evaluating the quantitative (PSNR, SSIM, *etc.*) or qualitative (visual subjective quality) differences of enhanced images *w.r.t.* the ground truths. Some earlier classical datasets include Set5 [49], Set14 [50], and LIVE1 [51]. The numbers of their images are small. Subsequent datasets come

with more diverse scene content, such as BSD500 [52] and Urban100 [53]. The popularity of deep learning methods has increased demand for training and testing data. Therefore, many newer and larger datasets are presented for image and video restoration, such as DIV2K [54] and MANGA109 [55] for image super-resolution, PolyU [56] and Darmstadt [57] for denoising, RawInDark [58] and LOL dataset [59] for low light enhancement, HazeRD [60], OHAZE [61] and IHAZE [62] for dehazing, Rain100L/H [17] and Rain800 [63] for rain streak removal, and RAINDROP [64] for raindrop removal. However, these datasets provide no integration with subsequent high-level tasks.

A few works [65–67] make preliminary attempts for event/action understanding, video summarization, or face recognition in unconstrained and potentially degraded environments. The following datasets are collected by aerial vehicles, including VIRAT Video Dataset [68] for event recognition, UAV123 [69] for UAV tracking, and a multi-purpose dataset [70]. In [71], an unconstrained Face Detection Dataset (UFDD) is proposed for face detection in adverse condition including weather-based degradations, motion blur, focus blur and several others, containing a total of 6,425 images with 10,897 face annotations. However, few works specifically consider the impacts of image enhancement and object detection/recognition jointly. Prior to this UG²⁺ effort, a number of latest works have taken the first stabs. A large-scale hazy image dataset and a comprehensive study – REAListic Single Image DEhazing (RESIDE) [72] – including paired synthetic data and unpaired real data is proposed to thoroughly examine visual reconstruction and vision recognition in hazy images. In [73], an Exclusively Dark (ExDARK) dataset is proposed with a collection of 7,363 images captured from very low-light environments with 12 object classes annotated on both image class level and local object bounding boxes. In [74], the authors present a new large-scale benchmark called RESIDE and a comprehensive study and evaluation of existing single image deraining algorithms, ranging from full-reference metrics, to no-reference metrics, to subjective evaluation and the novel task-driven evaluation. Those datasets and studies shed new light on the comparisons and limitations of state-of-the-art algorithms, and suggest promising future directions. In this work, we follow the footsteps of predecessors to advance the fields by proposing new benchmarks.

B. Poor Visibility Enhancement

There are numerous algorithms aiming to enhance visibility of the degraded imagery, such as image and video denoising/inpainting [75–79], deblurring [80, 81, 162? ?], super-resolution [82–85] and interpolation [86]. Here we focus on dehazing, low-light condition, and deraining, as in the UG²⁺ Track 2 scope.

Dehazing. Dehazing methods proposed in an early stage rely on the exploitation of natural image priors and depth statistics, *e.g.* locally constant constraints and decorrelation of the transmission [87], dark channel prior [9], color attenuation prior [88], nonlocal prior [89]. In [90, 91], Retinex theory is utilized to approximate the spectral properties of

object surfaces by the ratio of the reflected light. Recently, Convolutional Neural Network (CNN)-based methods bring in the new prosperity for dehazing. Several methods [6, 7] rely on various CNNs to learn the transmission fully from data. Beyond estimating the haze related variables separately, successive works make their efforts to estimate them in a unified way. In [92, 93], the authors use a factorial Markov random field that integrates the estimation of transmission and atmosphere light. Some researchers focus on the more challenging night-time dehazing problem [94, 95]. In addition to image dehazing, AOD-Net [8, 96] considers the joint interplay effect of dehazing and object detection in a unified framework. The idea is further applied to video dehazing by extending the model into a light-weight video hazing framework [97]. In another recent work [98], the semantic prior is also injected to facilitate video dehazing.

Low Light Enhancement. All low-light enhancement methods can be categorized into three ways: hand-crafted methods, Retinex theory-based methods and data-driven methods. Hand-crafted methods explore and apply various image priors to single image low-light enhancement, *e.g.* histogram equalization [99, 100]. Some methods [101, 102] regard the inverted low-light images as hazy images, and enhance the visibility by applying dehazing. The retinex theory-based method [103] is designed to transform the signal components, reflectance and illumination, differently to simultaneously suppress the noise and preserve high-frequency details. Different ways [104, 105] are used to decompose the signal and diverse priors [106–109] are applied to realize better light adjustment and noise suppression. Li *et al.* [20] further extended the traditional Retinex model to a robust one with an explicit noise term, and made the first attempt to estimate a noise map out of that model via an alternating direction minimization algorithm. A successive work [21] develops a fast sequential algorithm. Learning based low-light image enhancement methods [22–24] have also been studied, where low-light images used for training are synthesized by applying random Gamma transformation on natural normal light images. Some recent works aim to build paired training data from real scenes. In [58], Chen *et al.* introduced a See-in-the-Dark (SID) dataset of short-exposure low-light raw images with corresponding long-exposure reference raw images. Cai *et al.* [110] built a dataset of under/over-contrast and normal-contrast encoded image pairs, in which the reference normal-contrast images are generated by Multi-Exposure image Fusion or High Dynamic Range algorithms. Recently, Jiang *et al.* [111] proposed for the first time an unsupervised generative adversarial network, that can be trained without low/normal-light image pairs, yet generalizing nicely and flexibly on various real-world images. **Deraining.** Single image deraining is a highly ill-posed problem. To address it, many models and priors are used to perform signal separation and texture classification. These models include sparse coding [112], generalized low rank model [10], nonlocal mean filter [113], discriminative sparse coding [11], Gaussian mixture model [12], rain direction prior [13], transformed low rank model [14]. The presence of deep learning has promoted the development of single image deraining. In [15, 16], deep networks take the image detail

TABLE I
SUB-CHALLENGE 2.1: IMAGE AND OBJECT STATISTICS OF THE
TRAINING/VALIDATION, AND THE HELD-OUT TEST SETS.

	#Images	#Bounding Boxes
Training/Validation	4,310	41,113
Test (held-out)	2,987	24,201

TABLE II
SUB-CHALLENGE 2.1: CLASS STATISTICS OF THE TRAINING/VALIDATION,
AND THE HELD-OUT TEST SETS.

Categories	<i>Car</i>	<i>Person</i>	<i>Bus</i>	<i>Bicycle</i>	<i>Motorcycle</i>
RTTS	25,317	11,366	2,590	698	1,232
Test (held-out)	18,074	1,562	536	225	3,804

layer as their input. Yang *et al.* [17] proposed a deep joint rain detection and removal method to remove heavy rain streaks and accumulation. In [13], a novel density-aware multi-stream densely connected CNN is proposed for joint rain density estimation and removal. Video deraining can additionally make use of the temporal context and motion information. The early works formulate rain streaks with more flexible and intrinsic characteristics, including rain modeling [10, 114–124]. The presence of learning-based method [125–131], with improved modeling capacity, brings new progress. The emergence of deep learning-based methods push performance of video deraining to a new level. Chen *et al.* [132] integrated superpixel segmentation alignment, and consistency among these segments and CNN-based detail compensation network into a unified framework. Liu *et al.* [133] presented a recurrent network integrating rain degradation classification, deraining and background reconstruction.

C. Visual Recognition under Adverse Conditions

A real-world visual detection/recognition system needs to handle a complex mixture of both low-quality and high-quality images. It is commonly observed that, mild degradations, *e.g.* small noises, scaling with small factors, lead to almost no change of recognition performance. However, once the degradation level passes a certain threshold, there will be an unneglected or even very significant effect on system performance. In [134], Torralba *et al.* showed that, there will be a significant performance drop in object and scene recognition when the image resolution is reduced to 32×32 pixels. In [135], the boundary where the face recognition performance is largely degraded is 16×16 pixels. Karahan *et al.* [136] found Gaussian noise with its standard deviation ranging from 10 to 20 will cause a rapid performance decline. In [137], more impacts of contrast, brightness, sharpness, and out-of-focus on face recognition are analyzed.

In the era of deep learning, some methods [138–140] attempt to first enhance the input image and then forward the output into a classifier. However, this separate consideration of enhancement may not benefit the successive recognition task, because the first stage may incur artifacts which will damage the second stage recognition. In [135, 141], class-specific features are extracted as a prior to be incorporated into the restoration model. In [39], Zhang *et al.* developed a joint image restoration and recognition method based on sparse

representation prior, which constrains the identity of the test image and guides better reconstruction and recognition. In [8], Li *et al.* considered dehazing and object detection jointly. These two stage joint optimization methods achieve better performance than previous one-stage methods. In [40, 142], the joint optimization pipeline for low-resolution recognition is examined. In [41, 42], Liu *et al.* discussed the impact of denoising for semantic segmentation and advocated their mutual optimization. Lately, in [143], the algorithmic impact of enhancement algorithms for both visual quality and automatic object recognition is thoroughly examined, on a real image set with highly compound degradations. In our work, we take a further step to consider the joint enhancement and detection in bad weather environments. Three large-scale datasets are collected to inspire new ideas and novel methods in the related fields.

III. INTRODUCTION OF UG²⁺ TRACK 2 DATASETS

A. (Semi-)Supervised Object Detection in the Haze

In Sub-challenge 2.1, we use the 4,322 annotated real-world hazy images of the RESIDE RTTS set [72] as the training and/or validation sets (the split is up to the participants). Five categories of objects (car, bus, bicycle, motorcycle, pedestrian) are labeled with tight bounding boxes. We provide another 4,807 unannotated real-world hazy images collected from the same traffic camera sources, for the possible usage of semi-supervised training. The participants can optionally use pre-trained models (*e.g.*, on ImageNet or COCO), or external data. But if any pre-trained model, self-synthesized or self-collected data is used, that must be explicitly mentioned in their submissions, and the participants must ensure all their used data to be public available at the time of challenge submission, for reproducibility purposes.

There is a held-out test set of 2,987 real-world hazy images, collected from the same sources, with the same classes of objected annotated. Fig. 1 shows the basic statistics of the RTTS set and the held-out set. The held-out test set has a similar distribution of number of bounding boxes per image, bounding box size and relative scale of bounding boxes to input images compared to the RTTS set, but has relatively larger image size. Samples from RTTS set and held-out set can be found in Fig. 2 and Fig. 3.

B. (Semi-)Supervised Face Detection in the Low Light Condition

In Sub-challenge 2.2, we use our self-curated DARK FACE dataset. It is composed of 10,000 images (6,000 for training and validation, and 4,000 for testing) taken in under-exposure condition where human faces are annotated by human with bounding boxes; and 9,000 images taken with the same equipment in the similar environment without human annotations. Additionally, we provide a unique set of 789 paired low-light/normal-light images captured in controllable real lighting conditions (but unnecessarily containing faces), which can be optionally used as parts of the training data. The training and

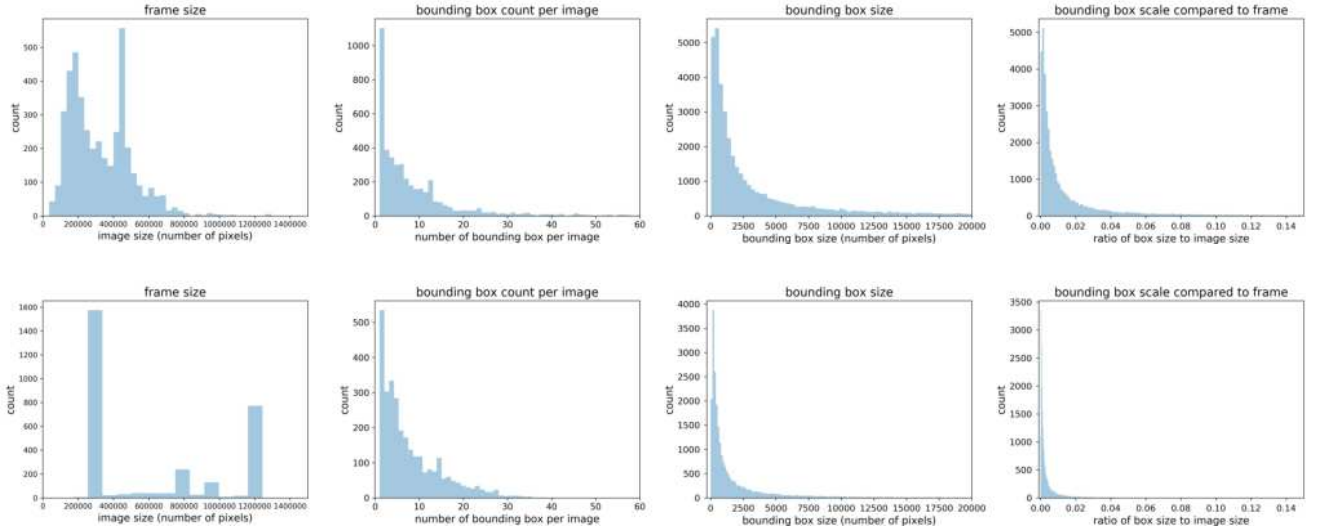


Fig. 1. Sub-challenge 2.1: Basic statistics on the training/validation set (the top row) and the held out test set (the bottom row). The first column shows the image size distribution (number of pixels per image), The second column the bounding box count distribution (number of bounding boxes per image), the third column the bounding box size distribution (number of pixels per bounding box), and the last column the ratios of bounding box size compared to frame size.



Fig. 2. Sub-challenge 2.1: Examples of images in training/validation set (*i.e.*, RESIDE RTTS [72]).

TABLE III
SUB-CHALLENGE 2.2: COMPARISON OF LOW-LIGHT IMAGE UNDERSTANDING DATASETS.

Dataset	Training		Testing	
	#Image	#Face	#Image	#Face
ExDark	400	-	209	-
UFDD	-	-	612	-
DarkFace	6,000	43,849	4,000	37,711

evaluation set includes 43,849 annotated faces and the held-out test set includes 37,711 annotated faces. Table III presents a summary of the dataset.

Collection and annotation. This collection consists of images recorded from Digital Single Lens Reflexes, specifically Sony $\alpha 6000$ and $\alpha 7$ E-mount cameras with different capturing parameters on several busy streets around Beijing, where faces of various scales and poses are captured. The images in this



Fig. 3. Sub-challenge 2.1: Examples of images in the held-out test set.

collection are open source content tagged with an Attribution-NonCommercial-NoDerivatives 4.0 International license¹. The resolution of these images is 1080×720 (down-sampled from $6K \times 4K$). After filtering out those without sufficient information (lacking faces, too dark to see anything, *etc.*), we select 10,000 images for human annotation. The bounding boxes are labeled for all the recognizable faces in our collection. We make the bounding boxes tightly around the forehead,

¹<https://www.jet.org.za/clearinghouse/projects/printed/resources/creative-commons-licence>

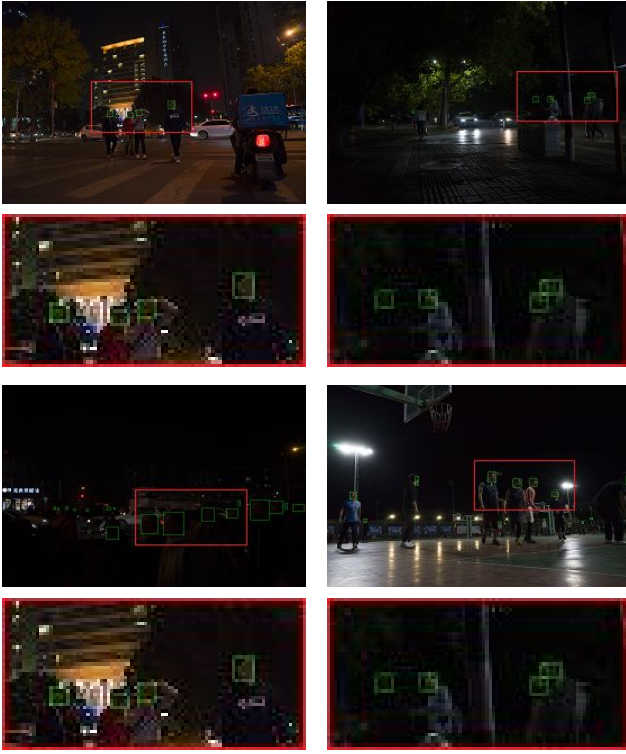


Fig. 4. Sub-challenge 2.2: DARK FACE has a high degree of variability in scale, pose, occlusion, appearance and illumination. The face regions in the red boxes are zoomed-in for better viewing.

TABLE IV
SUB-CHALLENGE 2.3: OBJECT STATISTICS IN THE HELD-OUT TEST SET.

Categories	Car	Person	Bus	Bicycle	Motorcycle
Test Set	7332	1135	613	268	968

chin, and cheek, using the LabelImg Toolbox². If a face is occluded, we only label the exposed skin region. If most of a face is occluded, we ignore it. For this collection, we observe commonly seen degradations in addition to under-exposure, such as intensive noise. The face number and resolution range distribution are displayed in Fig 5. Each annotated image contains 1-34 human faces. The face resolutions in these images range from 1×2 to 335×296 . The resolution of most faces in our dataset is below 300 pixel² and the face number mostly falls into the range [1, 20].

C. Zero-Shot Object Detection with Raindrop Occlusions

In Sub-challenge 2.3, we release 1,010 pairs of realistic raindrop images and corresponding clean ground-truths, collected from the real scenes as described in [64], as the training and/or validation sets. Our held-out test set contains 2,495 real rainy images from high-resolution driving videos. As shown in Fig. 6, all images are contaminated by raindrops on camera lens. They are captured in diverse real traffic locations and scenes during multiple drives. We label bounding boxes for selected traffic objects: car, person, bus, bicycle, and motorcycle, which commonly appear on the roads of all

images. Most images are of 1920×990 resolution, with a few exceptions of 4023×3024 resolution. The participants are free to use pre-trained models (e.g., ImageNet or COCO) or external data. But if any pre-trained model, self-synthesized or self-collected data is used, that must be explicitly mentioned in their submissions, and the participants must ensure their used data to be public available at the time of challenge submission, for reproducibility purposes.

D. Ranking Criterion

The ranking criteria will be the Mean average precision (mAP) on each held-out test set, with a default Intersection-of-Union (IoU) threshold as 0.5. If the ratio of the intersection of a detected region with an annotated object is greater than 0.5, a score of 1 is assigned to the detected region, and 0 otherwise. When mAPs with IoU as 0.5 are equal, the mAPs with higher IoUs (0.6, 0.7, 0.8) will be compared sequentially.

IV. BASELINE RESULTS AND ANALYSIS

For all three sub-challenges, we report results by **cascading off-the-shelf enhancement methods and popular pre-trained detectors**. There has been no joint training performed, hence the baseline numbers are in no way very competitive. We expect to see much performance boosts over the baselines from the competition participants.

A. Sub-challenge 2.1 Baseline Results

1) *Baseline Composition*: We test four state-of-the-art object detectors: (1) **Mask R-CNN**³ [144]; (2) **RetinaNet**⁴ [146]; and (3) **YOLO-V3**⁵ [147]; (4) Feature Pyramid Network⁶ (**FPN**) [148]. We also try three state-of-the-art dehazing approaches: (a) **AOD-Net**⁷ [8]; (b) Multi-Scale Convolutional Neural Network (**MSCNN**)⁸ [7]; (c) Densely Connected Pyramid Dehazing Network (**DCPDN**)⁹ [149]. All dehazing models adopt officially released versions.

2) *Results and Analysis*: We evaluate the object detection performance on the original hazy images of RESIDE RTTS set using Mask R-CNN. The detectors are pretrained on Microsoft COCO, a large-scale object detection, segmentation, and captioning dataset. The detailed detection performance on the five objects can be found in Table V. Results show that without preprocessing or dehazing, the object detectors pretrained on clean images fail to predict a large amount of objects in the hazy image. The overall detection performance has a mAP of only 41.83% using Mask R-CNN and 42.54% using YOLO-V3. Among all the five object categories, person has the highest detection AP, while bus has the lowest AP.

We also compare the validation and test set performance in Table. V. One possible reason for the performance gap between validation and test sets is that the bounding box size

²<https://github.com/tzutalin/labelImg>

³https://github.com/matterport/Mask_RCNN

⁴<https://github.com/fizyr/keras-retinanet>

⁵<https://github.com/ayoozhkathuria/pytorch-yolo-v3>

⁶https://github.com/DetectionTeamUCAS/FPN_Tensorflow

⁷<https://github.com/Boyiilee/AOD-Net>

⁸<https://github.com/rwenqi/Multi-scale-CNN-Dehazing>

⁹<https://github.com/hezhangsprinter/DCPDN>

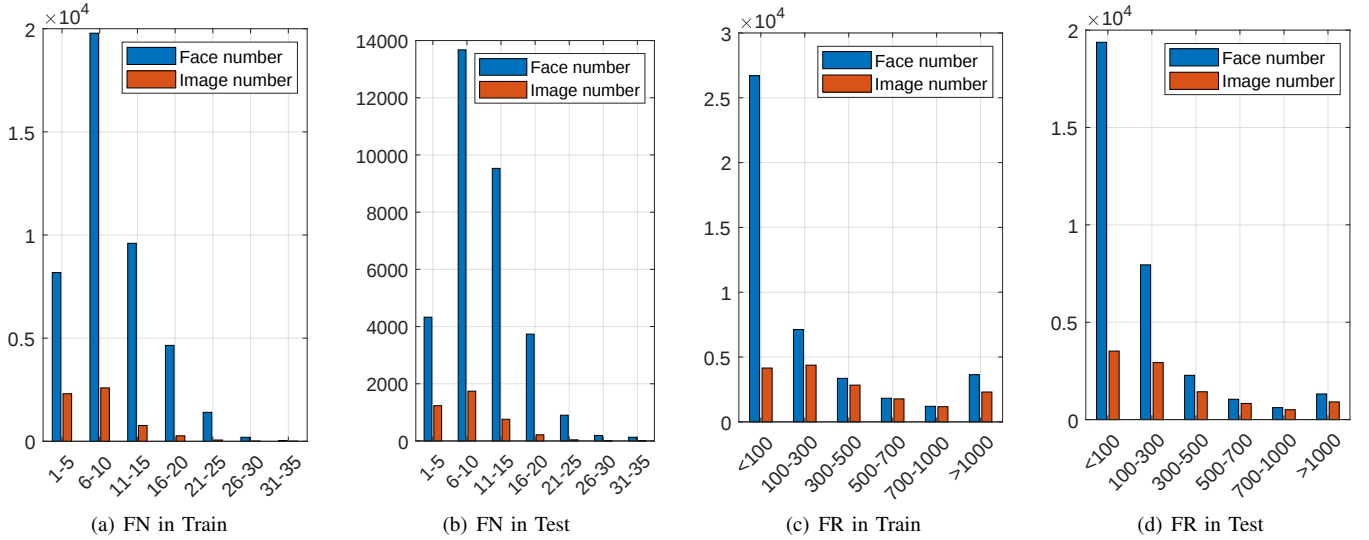


Fig. 5. Sub-challenge 2.2: Face resolution (FR) and face number (FN) distribution in DARK FACE collections. Image number denotes the number of images belonging to a certain category. Face number denotes the summation number of faces belonging to a certain category.



Fig. 6. Sub-challenge 2.3: Example images from the held-out test set.

of the latter is smaller compared to the former, as showed in Fig. 1 as well as visualized in Fig. 7.

Besides, we analyze the difference between the synthetic haze/rain images and those in real applications. The haze image is generated from the model:

$$I = Jt + A(1 - t), \quad (1)$$

where I is the observed hazy image, J is the scene reliance to be recovered. A denotes the global atmospheric light, and t is the transmission matrix assumed correlated with the scene depth. For synthesis hazy images as shown at the top panel of Fig. 8, t is strictly inferred from the results of depth estimation. Therefore, the haze is distributed more homogeneously among different regions and objects, as existing depth estimation techniques are not adaptive enough to accurately estimate the fine-grained depth of each object. Comparatively, the haze layers of real hazy images as shown in the bottom panel of Fig. 8 might be uncorrelated to the scene depth or have more variations for objects with various depths in single hazy

images. Besides, the scattering and refraction of light under real hazy condition is different to that in clear environment. The glow effects in front of the vehicles in haze (e.g. last examples) makes the vehicle recognition more difficult using pretrained object detection algorithms.

3) *Effect of Dehazing*: We further evaluate the current state-of-the-art dehazing approaches on hazy dataset, with pre-trained detectors subsequently applied without tuning or adaptation. Fig. 7 shows two examples that dehazing algorithms can improve not only the visual quality of the images but also the detection accuracies. More detection results are included in Table. V. Detection mAPs of dehazed images using DCPDN and MSCNN approaches are 1% higher on average compared to those of hazy images. Eventually, the choice of pre-trained detectors seem to also matter here: Mask R-CNN outperforms the other two detectors on both validation and test sets, before and after dehazing.

Furthermore, as reported in [149] and Table VI, DCPDN has the best SSIM scores while MSCNN has the worst visual quality. However, the detection performance of MSCNN is much better than that of DCPDN.

B. Sub-challenge 2.2 Baseline Results

1) *Baseline Composition*: We test four state-of-the-art deep face detectors: (1) Dual Shot Face Detector (**DSFD**) [150]¹⁰; (2) **Pyramidbox** [151]¹¹; (3) Single Stage Headless Face Detector (**SSH**) [152]¹²; (4) **Faster RCNN** [153]¹³.

We also include seven state-of-the-art algorithms for light/contrast enhancement: (a) Bio-Inspired Multi-Exposure Fusion (**BIMEF**) [154]¹⁴; (b) **Dehazing** [155]¹⁴; (c) Low-light Image Enhancement (**LIME**) [108]¹⁵; (d) **MF** [107]¹⁴; (e)

¹⁰<https://github.com/TencentYoutuResearch/FaceDetection-DSFD>

¹¹<https://github.com/EricZgw/PyramidBox>

¹²<https://github.com/mahyarnajibi/SSH.git>

¹³<https://github.com/playerkk/face-py-faster-rcnn>

¹⁴<https://github.com/baidut/BIMEF>

¹⁵<https://sites.google.com/view/xjguo/lime>



Fig. 7. Examples of object detection of hazy images and dehazed images on RESIDE RTTS set. The first column displays the ground truth bounding boxes on hazy images, the second column displays detected bounding box on hazy image using pretrained Mask R-CNN, the right three columns display Mask R-CNN detected bounding boxed on dehazed images using AOD-Net, MSCNN, DCPDN correspondingly.

TABLE V
DETECTION RESULTS (MAP) ON THE RTTS (TRAIN/VALIDATION DATASET) AND HELD-OUT TEST SETS.

mAP			hazy	AOD-Net [8]	DCPDN [7]	MSCNN [149]
validation	* RetinaNet [146]	Person	55.85	54.93	56.70	58.07
		Car	41.19	37.61	42.68	42.77
		Bicycle	39.61	37.80	36.96	38.16
		Motorcycle	27.37	23.31	29.18	29.01
		Bus	16.88	15.70	16.34	18.34
		mAP	36.18	33.87	36.37	37.27
	* Mask R-CNN [144]	Person	67.52	66.71	67.18	69.23
		Car	48.93	47.76	52.37	51.93
		Bicycle	40.81	39.66	40.40	40.42
		Motorcycle	33.78	26.71	34.58	31.38
		Bus	18.11	16.91	18.25	18.42
		mAP	41.83	39.55	42.56	42.28
	* YOLO-V3 [147]	Person	60.81	60.21	60.42	61.56
		Car	47.84	47.32	48.17	49.75
		Bicycle	41.03	42.22	40.18	42.01
		Motorcycle	39.29	37.55	38.17	41.11
		Bus	23.71	20.91	23.35	23.15
		mAP	42.54	41.64	42.06	43.52
	◇ FPN [148]	Person	51.85	52.35	51.04	54.50
		Car	37.48	36.05	37.19	38.88
		Bicycle	35.31	35.93	32.57	37.01
		Motorcycle	23.65	21.07	22.97	23.86
		Bus	12.95	13.68	12.07	15.83
		mAP	32.25	31.82	31.17	34.02
test	RetinaNet	Person	17.64	18.23	16.65	19.34
		Car	31.41	29.30	33.31	32.97
		Bicycle	0.42	0.84	0.38	0.75
		Motorcycle	1.69	1.37	1.93	2.03
		Bus	12.77	13.70	12.07	15.82
		mAP	12.79	12.69	12.87	14.18
	Mask R-CNN	Person	25.60	26.63	24.59	27.94
		Car	39.31	39.71	42.76	42.57
		Bicycle	0.64	0.52	0.22	0.37
		Motorcycle	3.37	2.81	2.83	2.99
		Bus	15.66	15.41	16.69	16.55
		mAP	16.92	17.02	17.42	18.09
	YOLO-V3	Person	20.64	21.41	21.42	22.11
		Car	34.68	33.90	34.52	35.93
		Bicycle	0.50	0.38	0.98	0.57
		Motorcycle	4.26	4.10	4.72	5.27
		Bus	13.55	14.35	13.75	15.04
		mAP	14.69	14.83	15.08	15.78
	FPN	Person	12.65	12.57	11.13	14.19
		Car	30.54	31.24	27.81	32.68
		Bicycle	1.91	0.39	1.12	0.97
		Motorcycle	2.25	1.7	1.96	1.89
		Bus	6.08	7.93	7.39	8.31
		mAP	10.69	10.77	9.88	11.61

* RetinaNet, Mask R-CNN and YOLO-V3 are pretrained on Microsoft COCO dataset.

◇ FPN using ResNet-101 backbone is pretrained on the PASCAL Visual Object Classes (VOC) dataset.

Multi-Scale Retinex (MSR) [105]¹⁴; (f) Joint Enhancement and Denoising (JED) [21]¹⁶; (g) RetinexNet [59]¹⁷.

¹⁶<https://github.com/tonghelen/JED-Method>



Fig. 8. Compared with synthetic hazy images in OTS dataset (at the top panel), the haze layers of real hazy images (at the bottom panel) in RTTS dataset might be uncorrelated to the scene depth and have more variations for objects with various depths in single hazy images.

TABLE VI
COMPARISON OF HUMAN AND MACHINE VISION QUALITY DIFFERENT METHODS ACHIEVE.

Metric	Dataset	Baseline	MSCNN	AOD-Net	DCPDN
SSIM	TestA [140]	-	0.8203	0.8842	0.956
SSIM	TestB [140]	-	0.7724	0.8325	0.8746
MAP	UG2.1-Test	RetinaNet	14.18	12.69	12.87
		Mask R-CNN	18.09	17.02	17.42
		YOLO-V3	15.78	14.83	15.08
		FPN	11.61	10.77	9.88

2) *Results and Analysis*: Fig. 14 (a) depicts the precision-recall curves of the original face detection methods, without enhancement. The baseline methods are trained on WIDER FACE [156]¹⁸, a large dataset with large scale variations in diversified factors and conditions. The results demonstrate that without proper pre-processing or adaptation, the state-of-the-art methods cannot achieve desirable detection rates on DARK FACE. Result examples are illustrated in Fig. 12. The evidences may imply that previous face datasets, though covering variations in poses, appearances, scale, *et al.*, are still insufficient to capture the facial features in the highly under-exposed condition.

We have showed some failure cases of Sub-challenge 2.2. In these results, the faces are falsely detected (false positives and false negatives) due to heavy degradation, small scale, pose variation, occlusion, *etc.* As show in Fig. 9, due to the above mentioned reasons, there are a certain amount of false negative samples caused by heavy degradation, small scale, pose variation, occlusion at the top four panels, respectively, and false positive samples at the bottom panel. DSFD is used as the baseline method. For better visibility, the results shown here are processed by LIME.

3) *Effect of Enhancement*: We next use the enhancement algorithms to pre-process the annotated dataset and then apply the above two pre-trained face detection methods to the processed data. While the visual quality of the enhanced images is better, as expected, the detectors do perform better. As shown in Fig. 14 (b) and (c), in most instances, the precision of the detectors notably increases compared to that of the data without enhancement. Various existing enhancement methods seem to result in similar improvements here.

Despite being encouraging to see, the overall performance of the detectors still drops a lot compared to normal-light

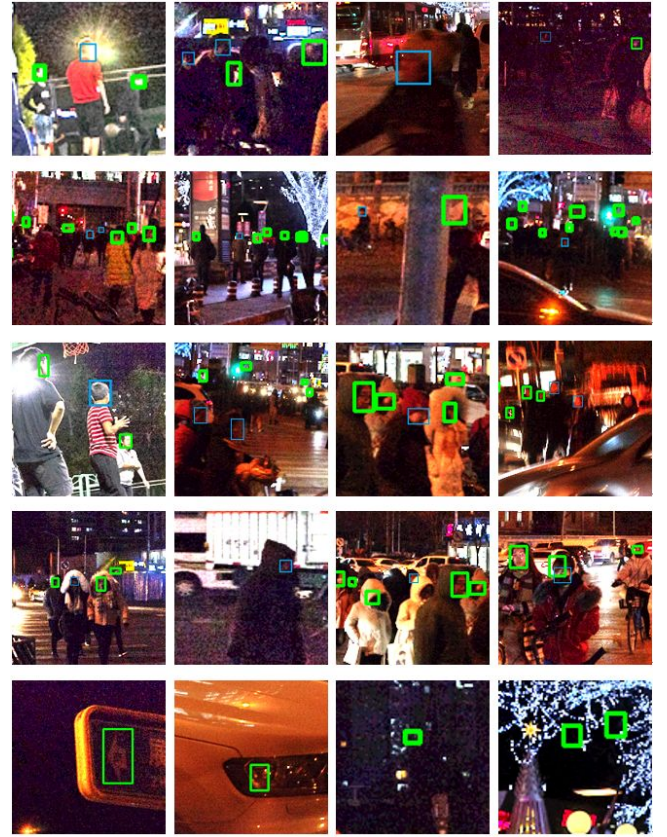


Fig. 9. Failure case analysis even with the model being trained with the proposed training set. Green Box: detection results by the baseline DSFD. Blue boxes: ground truths that are not detected.

datasets. The simple cascade of low light enhancement and face detectors leave much improvement room open.

C. Sub-challenge 2.3 Baseline Results

1) *Baseline Composition*: We use four state-of-the-art object detection models: (1) Faster R-CNN (**FRCNN**) [157]; (2) **YOLO-V3** [147]; (3) **SSD-512** [158]; and (4) **RetinaNet** [146].

We employ five state-of-the-art deep learning-based deraining algorithms: (a) JOint Rain DEtection and Removal¹⁹ (**JORDER**) [17]; (b) Deep Detail Network²⁰ (**DDN**) [15]; (c) Conditional Generative Adversarial Network²¹ (**CGAN**) [63]; (d) Density-aware Image De-raining method using a Multistream Dense Network²² (**DID-MDN**) [13]; and (e) **DeRaindrop**²³ [64]. For fair comparisons, we re-train all deraining algorithms using the same provided training set.

2) *Results and Analysis*: Table VII shows mAP results comparisons for different deraining algorithms using different detection models on the held-out test set. Unfortunately, we find that almost all existing deraining algorithms deteriorate the objects detection performance compared to directly using

¹⁷<https://github.com/weichen582/RetinexNet>

¹⁸<http://shuoyang1213.me/WIDERFACE/>

¹⁹http://www.icst.pku.edu.cn/struct/Projects/joint_rain_removal.html

²⁰https://github.com/XMU-smartdsp/Removing_Rain

²¹https://github.com/TrinhQuocNguyen/Edited_Original_IDCGAN

²²<https://github.com/hezhangsprinter/DID-MDN>

²³<https://github.com/rui1996/DeRaindrop>

the rainy images for YOLO-V3, SSD-512, and RetinaNet (The only exception is the detection results by FRCNN). This could be due to those deraining algorithms are not trained towards the end goal of object detection, they are unnecessary to help this goal, and the deraining process itself might have lost discriminative, semantically meaningful true information, and thus hampers the detection performance. In addition, Table VII shows that YOLO-V3 achieves the best detection performance, independent of deraining algorithms applied. We attribute this to the small objects in a relative long distance from the camera in the test set since YOLO-V3 is known to improve small object detection based on multi-scale prediction structure.

V. COMPETITION RESULTS: OVERVIEW & ANALYSIS

The UG²+ Challenge (Track 2) in conjunction with CVPR 2019 attracted large deals of attention and participation. More than 260 teams registered; among them, 82 teams finished the dry-run and submitted their final results successfully. Eventually, 6 teams were selected as winners (including a winner and a runner-up, for each sub-challenge).

In the following, we review a part of results from those participation teams who volunteer to disclose their technical details. The full leaderboards can be found at the website ²⁴.

A. Sub-challenge 2.1: Competition Results and Analysis

A total of seven teams were able to outperform our best baseline numbers (mAP 18.09). The winner and runner-up teams, *HRI_DET* and *superlab403*, achieve record-high mAP results of 52.71 and 49.22, respectively. All teams used deep learning solutions. In addition to using most sophisticated networks, several interesting observations could be concluded: i) while many teams went with the dehazing-detection cascade idea (like our baselines, but usually jointly trained), the top-2 winners used end-to-end trained/adapted detection models on the hazy training set, without an (implicit) dehazing module; ii) the utilization of unlabeled data seems to open up much potential, and we believe it should be paid more attention to in the future; and iii) multi-scale testing and ensembling contribute to many performance gains.

As the winner team, *HRI_DET* used Faster R-CNN, with the ImageNet-pretrained backbone of ResNeXt-101 [159] and Feature Pyramid Network (FPN) [148]. The Faster R-CNN was then tuned on the mixed dataset of MS COCO, PASCAL VOC and KITTI, with the common data augmentations. To further boost the performance, the team delved deep into the provided unlabeled hazy set, and adopted semi-supervised learning by using the unlabeled data to train the feature extractor with a reconstruction loss. The team used a batch size of 8 and trained the network using 8 Tesla M40 GPUs for 30,000 iterations with an initial learning rate of 0.005. They also found it helpful to apply stochastic weight averaging (SWA) [160] to aggregate several checkpoint models in one training pass, and further to ensemble multiple models (by averaging model weights) obtained from different training passes (e.g. with different learning rate schedulers). During

inference, a three-scale testing is performed by resizing images to 1333×1000 , 1000×750 and 2100×600 . The third scale is applied to a closer view of the image, by cropping the foreground region defined as the bounding box of all high-confidence predictions with the first scale.

The runner-up team *superlab403* chose the Cascade R-CNN [161] baseline, and also replaced the original ResNet-101 backbone with ResNeXt-101. The team analyzed the distribution of target aspect ratio from the training set, and selected four new anchor ratios (0.8, 1.7, 2.6, 3.7) by k -means. The team also did a (well-appreciated) label cleaning effort. Data augmentations such as blur, illumination change and color perturbations were adopted in training. The model was trained with an SGD optimizer; the initial learning rate was set as 0.0025, then being decayed by a factor of 0.1 at epochs 8, 11, 21 and 41 (total training epoch number 50).

Other teams have each developed their interesting solutions. For example, the *Mt. Star* team (ranked No. 3, mAP 31.24) adopted a sequential cascade of the dehazing model (DehazeNet [6]) and the detection model (Faster-RCNN), each first pre-trained on their own and then jointly tuned end-to-end on the training set. Multi-scale testing was adopted. The *ilab* team (ranked No. 6, mAP 19.15) also referred to the dehazing-detection cascade idea, but using DeblurGAN [162] (re-trained on haze data) for the dehazing model and Yolo-V2 [147] for the detection model, with a content loss.

B. Sub-challenge 2.2: Competition Results and Analysis

A total of three teams were able to outperform our best baseline numbers (mAP 39.30). The winner and runner-up teams, *CAS-Newcastle* and *CAS_NEU*, achieve high mAP results of 62.45 and 61.84, respectively. Similarly to Sub-challenge 2.1, all teams used deep learning solutions; yet interestingly, the most successful solutions are based on enhancement-detection cascades, showing a different trend with Sub-challenge 2.1.

The *CAS-Newcastle* and *CAS_NEU* team adopted cascades of low-light enhancement (MSRCR [105]) and detection (Selective Refinement Network [145] / RetinaNet [146]) models, where the detection models were directly trained on the enhancement models' preprocessed outputs. The *SCUT-CVC* team (ranked No. 7, mAP 35.18) found tone mapping [163] to be an impressively effective pre-processing, on top of which they tuned two DSFD detectors (with VGG-16 and ResNet-152 backbones), whose results were ensembled by late fusion. The *PHI-AI* team (ranked No. 7, mAP 29.95) adopted a U-Net [164] enhancer and a DSFD detector. The *tjfirst* team (ranked No. 12, mAP 26.50) referred to a more sophisticated enhancement module (first enhancing illumination by LIME [108], then super-resolving by DPSR [165], ended by denoising with BM3D [4]), followed by aggregating the DSFD-detection results on the original and enhanced images.

C. Sub-challenge 2.3: Competition Results and Analysis

Different from the first two sub-challenges, Sub-challenge 2.3 is substantially more difficult due to its "zero-shot" nature. Typical solutions that we see from the challenge teams include deraining + detection cascades; as well as ensembling

²⁴http://www.ug2challenge.org/leaderboard19_t2.html



Fig. 10. Samples of open-source paired rain image training set only including the rain streak and raindrop-related degradation.



Fig. 11. Real rain images captured in driving or surveillance cases have degradation of rain accumulation, blurring, reflectance and occlusions.

multiple pre-trained detectors (e.g., the CAS-Newcastle-TUM team). Unfortunately but not too surprisingly, none of the participation teams was able to outperform our baseline. That concurs with the conclusion drawn from the recent benchmark work [74]: “Perhaps surprisingly at the first glance, we find that almost all existing deraining algorithms will deteriorate the detection performance compared to directly using the rainy images...” “No existing deraining method seems to directly help detection. That may encourage the community to develop new robust algorithms to account for high-level vision problems on real-world rainy images. On the other hand, to realize the goal of robust detection in rain does not have to adopt a de-raining preprocessing; there are other domain adaptation type options...”.

In fact, our Sub-challenge 2.3 is more difficult and challenging. There is no training set that is close to the testing set provided, which fails all submitted methods. Therefore, the problem is closer to zero-shot and unsupervised learning problem. Existing open-source paired rain image training set, e.g. Rain800 [63] and raindrop dataset [64], usually consider only one kind of rain degradation, i.e. rain streak or raindrop. Based on a benchmark paper [72], the synthetic and captured paired images rely on three rain models. The generated rain images are not visually authentic and close to the real captured ones. Their background layers of these images are clear and the objects in these images have their appropriate sizes and locations as shown in Fig. 10.

The testing rain images in our Sub-challenge 2.3 are collected in real driving or surveillance scenarios. They come with other degradation such as rain accumulation, blurring, reflectance, occlusion *etc.* as shown in Fig. 11, which causes the domain shift problem and makes the related restoration and detection tasks harder.

Our challenging results further confirm that, the existing open-source paired datasets are poorly related when it comes to task purposes, e.g. object detection, in real scenario. Since driving and surveillance are representative of real application scenarios where deraining may be desired, this sub-challenge

is well-worth exploration. We are intended to extend this challenge to next year and look for better deraining algorithms to be proposed.

VI. DISCUSSIONS AND FUTURE DIRECTIONS

In this challenge, the submitted solutions and analysis results provide rich experiences, meaningful observations and insights, as well as potential future directions:

- Deep learning methods are preferred by all participants in their submitted solutions. Except that in Sub-challenge 2.2, some teams choose to apply hand-crafted low-light enhancement methods as pre-processing, other submitted methods are deep learning-based as they are flexible to be tuned to improve the performance based on the given task (dataset).
- In Sub-challenge 2.1, the top-2 winners make efforts in exploring the potential of unlabeled data, via semi-supervised learning with a reconstruction loss to guide the reconstruction of the full-picture from the intermediate feature. This strategy leads to performance improvement, which shows that it should be paid more attention to in future.
- For different tasks and techniques used to tackle the problems, the best choices of the framework might be different. In our challenge, the winner of Sub-challenge 2.1 uses a one-step detection scheme (without an implicit dehazing module) while the winner of Sub-challenge 2.2 takes the cascade of enhancement and detection.
- Some teams, e.g. *tjfirst*, report the effectiveness to apply a sequential enhancement process to remove mixed degradation, such as under-exposure, low-resolution blurring, and noise. It also shows a valuable path that is worth further exploration.
- Separate consideration of enhancement and detection might lead to deteriorated performance in detection. In Table V, the offline dehazing operation is first applied by AOD-Net, DCPDN, MSCNN. The results are taken as the input of object detection methods, i.e. RetinaNet, Mask R-CNN, YOLO-V3, FPN. Many combination groups of dehazing operation and object detection method obtain inferior results than directly applying object detection without any pre-enhancement, such as (AOD-Net, RetinaNet) for person category ($54.93 < 55.85$) and (AOD-Net, Mask R-CNN), (AOD-Net, Mask R-CNN), (AOD-Net, Mask R-CNN) for all categories ($2.81, 2.83, 2.99 < 3.37$). In Table VII, YOLO-V3, SSD-512 and RetinaNet generate worse object detection results if pre-deraining is applied.
- The existing results show the difficulty of three tasks we set. In Sub-challenge 2.1 and 2.2, the winners only achieve the results below 65 MAP, much inferior to the performance on datasets with degradation. In Sub-challenge 2.3, no participants achieve the results superior to the baseline results. The results show that, our datasets are challenging and there is still large room for further improvement.

VII. CONCLUSIONS

As concurred by most teams in the post-challenge feedbacks, it is widely agreed that the three sub-challenges in the UG²⁺ challenge 2019 Track 2 represent a very difficult, under-explored, yet high meaningful class of computer vision problems in practice. While some promising progress has been witnessed from the large volume of team participation²⁵, there remains large room to be improved. Through organizing this challenge, we expect to evoke a broader attention from the

²⁵ Taiheng Zhang is with the Department of Mechanical Engineering, Zhejiang University, Hangzhou 310027, China (email: thzhang@zju.edu.cn).

Qiaoyong Zhong, Di Xie and Shiliang Pu are with the Hikvision Research Institute, Hangzhou 310051, China (e-mail: zhongqiaoyong@hikvision.com; xiedi@hikvision.com; pushiliang.hri@hikvision.com).

Hao Jiang, Siyuan Yang, Yan Liu, Xiaochao Qu, Pengfei Wan are with the Mtlab, Meitu Inc., Beijing 100080, China (e-mail: jh1@meitu.com, ysy2@meitu.com, ly33@meitu.com, qxc@meitu.com, wpf@meitu.com).

Shuai Zheng, Minhui Zhong, Lingzhi He, Zhenfeng Zhu and Yao Zhao are with the Institute of Information Science, Beijing Jiaotong University, Beijing, 100044, China (e-mail: ericzhen1997@163.com, zylayee@163.com, 19112002@bjtu.edu.cn, zhfhzhu@bjtu.edu.cn, yzhao@bjtu.edu.cn).

Taiyi Su is with the Department of Computer Science and Technology, Tongji University, Shanghai, 201804, China (e-mail: tysu@tongji.edu.cn). Yandong Guo is with XPENGMOTORS, Beijing, China (e-mail: guoyd@xiaopeng.com).

Jinxu Liang, Tianyi Chen, Yuhui Quan and Yong Xu are with the School of Computer Science and Engineering at South China University of Technology, Guangzhou 510006, China (e-mail: csssherryliang@mail.scut.edu.cn; csttychen@mail.scut.edu.cn; csyhquan@scut.edu.cn; yxu@scut.edu.cn).

Jingwen Wang is with Tencent AI Lab, Shenzhen 518000, China (e-mail: jwongwang@tencent.com).

Shifeng Zhang, Chubing Zhuang and Zhen Lei are with the Institute of Automation of the Chinese Academy of Sciences, Beijing 100190, China (e-mail: shifeng.zhang@nlpr.ia.ac.cn; chubin.zhuang@nlpr.ia.ac.cn; zlei@nlpr.ia.ac.cn).

Jianning Chi, Huan Wang and Yixiu Liu are with the Northeastern University, Shenyang, 110819, China (e-mail: chijianing@mail.neu.edu.cn; 1870692@stu.neu.edu.cn; yixiu9713@foxmail.com).

Huang Jing, Guo Heng and Jianfei Yang are with the Nanyang Technological University, 639798, Singapore (e-mail: jhuang027@e.ntu.edu.sg; hgao007@e.ntu.edu.sg; yang0478@e.ntu.edu.sg).

Zhenyu Chen is with the Big Data Center, State Grid Corporation of China, Beijing, China and the China Electric Power Research Institute, Beijing 100031, China (e-mail: czy9907@gmail.com).

Dong Yi is with the Winsense Inc., Beijing 100080, China (e-mail: yidong@winsense.ai).

Cheng Chi, Kai Wang, Jiangang Yang, Likun Qin are with the University of Chinese Academy of Sciences, Beijing 100049, China (chicheng15@mails.ucas.ac.cn; wk2008cumt@163.com; yangjiangang18@mails.ucas.ac.cn; qinlikun19@ucas.ac.cn).

Liguo Zhou and Mingyue Feng are with the Department of Informatics, Technical University of Munich, Garching 85748, Germany (e-mail: liguo.zhou@tum.de; mingyue.feng@tum.de).

Chang Guo, Yongzhou Li and Huicai Zhong are with the Institute of Microelectronics of the Chinese Academy of Sciences, Beijing 100190, China (e-mail: guochang@ime.ac.cn; liyongzhou@ime.ac.cn; zhonghuicai@ime.ac.cn).

Xingyu Gao is with the Institute of Microelectronics of the Chinese Academy of Sciences, Beijing, China and the Beijing Jiaotong University, Beijing 100044, China (e-mail: gaopingyu@ime.ac.cn).

Stan Z. Li is with the Institute of Automation of the Chinese Academy of Sciences, Beijing and the Westlake University, Hangzhou 310024, China (e-mail: szli@nlpr.ia.ac.cn).

Zheming Zuo is with the Department of Computer Science, Durham University, Durham, United Kingdom (e-mail: zheming.zuo@durham.ac.uk).

Wenjuan Liao is with the Australian National University, Acton ACT 0200, Australia (e-mail: wenjuan.liao@anu.edu.au).

Ruizhe Liu is with the Sunway-AI Co., Ltd, Zhuhai, China and the Chinese Academy of Sciences R&D Center for Internet of Things, Wuxi 214200, China (e-mail: liuruizhe@ciotc.org).

Shizheng Wang is with the Institute of Microelectronics of Chinese Academy of Sciences, Beijing, China and the Chinese Academy of Sciences R&D Center for Internet of Things, Wuxi 214200, China (e-mail: shizheng.wang@foxmail.com).



Fig. 12. Sample face detection results of pretrained baseline on the original images of the proposed DARK FACE dataset. The face regions in the red boxes are zoomed-in for better viewing.



Fig. 13. Sample face detection results of pretrained baseline on the enhanced images of the proposed DARK FACE dataset. The face regions in the red boxes are zoomed-in for better viewing.

research community to address these challenges, which are barely covered by previous benchmarks. We look forward to making UG²⁺ a recurring event and also evolving/updating our problems and datasets every year.

REFERENCES

- [1] Y. Zhang, S. Kwong, and S. Wang, "Machine learning based video coding optimizations: A survey," *Information Sciences*, vol. 506, pp. 395 – 423, 2020.
- [2] P. Zhu, L. Wen, D. Du, X. Bian, H. Ling, Q. Hu, Q. Nie, H. Cheng, C. Liu, X. Liu et al., "VisDrone-DET2018: The vision meets drone object detection in image challenge results," in *Proc. European Conf. on Computer Vision (ECCV)*, 2018, pp. 437–468.
- [3] J. Yang, J. Wright, T. S. Huang, and Y. Ma, "Image super-resolution via sparse representation," *IEEE Trans. on Image Processing*, vol. 19, no. 11, pp. 2861–2873, Nov 2010.
- [4] K. Dabov, A. Foi, V. Katkovnik, and K. Egiazarian, "Image denoising by sparse 3-d transform-domain collaborative filtering," *IEEE Trans. on Image Processing*, vol. 16, no. 8, pp. 2080–2095, Aug 2007.
- [5] D. G. Lowe, "Distinctive image features from scale-invariant keypoints," *International Journal of Computer Vision*, vol. 60, pp. 91–110, 2004.
- [6] B. Cai, X. Xu, K. Jia, C. Qing, and D. Tao, "Dehazenet: An end-to-end system for single image haze removal," *IEEE Trans. on Image Processing*, vol. 25, no. 11, pp. 5187–5198, November 2016.
- [7] W. Ren, S. Liu, H. Zhang, J. Pan, X. Cao, and M.-H. Yang, "Single image dehazing via multi-scale convolutional neural networks," in *Proc. European Conf. on Computer Vision (ECCV)*, 2016, pp. 154–169.
- [8] B. Li, X. Peng, Z. Wang, J. Xu, and D. Feng, "AOD-Net: All-in-one dehazing network," in *Proc. IEEE Int'l Conf. Computer Vision*, 2017, pp. 4780–4788.
- [9] K. He, J. Sun, and X. Tang, "Single image haze removal using dark channel prior," *IEEE Trans. on Pattern Analysis and Machine Intelligence*, vol. 33, no. 12, pp. 2341–2353, Dec 2011.
- [10] Y.-L. Chen and C.-T. Hsu, "A generalized low-rank appearance model for spatio-temporally correlated rain streaks," in *Proc. IEEE Int'l Conf. on Computer Vision*, 2013, pp. 1968–1975.
- [11] Y. Luo, Y. Xu, and H. Ji, "Removing rain from a single image via discriminative sparse coding," in *Proc. IEEE Int'l Conf. on Computer Vision*, 2015, pp. 3397–3405.

TABLE VII
DETECTION RESULTS (MAP) ON THE HELD-OUT TEST SET.

	Rainy	JORDER [17]	DDN [15]	CGAN [63]	DID-MDN [13]	DeRaindrop [64]
FRCNN [157]	16.52	16.97	18.36	23.42	16.11	15.58
YOLO-V3 [147]	27.84	26.72	26.20	23.75	24.62	24.96
SSD-512 [158]	17.71	17.06	16.93	16.71	16.70	16.69
RetinaNet [146]	23.92	21.71	21.60	19.28	20.08	19.73

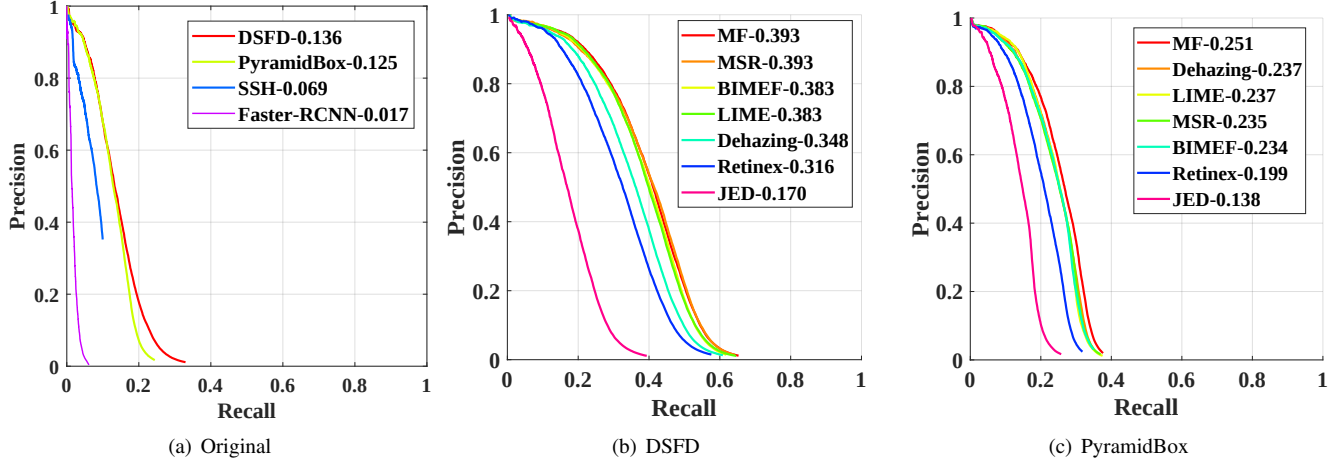


Fig. 14. Evaluation results of pretrained baseline on original and enhanced images of the proposed DARK FACE dataset.

- [12] Y. Li, R. T. Tan, X. Guo, J. Lu, and M. S. Brown, "Rain streak removal using layer priors," in *Proc. IEEE Int'l Conf. Computer Vision and Pattern Recognition*, 2016, pp. 2736-2744.
- [13] H. Zhang and V. M. Patel, "Density-aware single image de-raining using a multi-stream dense network," in *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, 2018, pp. 695-704.
- [14] Y. Chang, L. Yan, and S. Zhong, "Transformed low-rank model for line pattern noise removal," in *Proc. IEEE Int'l Conf. Computer Vision*, Oct 2017, pp. 1735-1743.
- [15] X. Fu, J. Huang, D. Zeng, Y. Huang, X. Ding, and J. Paisley, "Removing rain from single images via a deep detail network," in *Proc. IEEE Int'l Conf. Computer Vision and Pattern Recognition*, 2017, pp. 1715-1723.
- [16] X. Fu, J. Huang, X. Ding, Y. Liao, and J. Paisley, "Clearing the skies: A deep network architecture for single-image rain removal," *IEEE Trans. on Image Processing*, vol. 26, no. 6, pp. 2944-2956, June 2017.
- [17] W. Yang, R. T. Tan, J. Feng, J. Liu, Z. Guo, and S. Yan, "Deep joint rain detection and removal from a single image," in *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, 2017, pp. 1685-1694.
- [18] W. Yang, R. T. Tan, J. Feng, J. Liu, S. Yan and Z. Guo, "Joint Rain Detection and Removal from a Single Image with Contextualized Deep Networks," in *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2020.
- [19] R. Qian, R. T. Tan, W. Yang, J. Su and J. Liu, "Attentive Generative Adversarial Network for Raindrop Removal from A Single Image," in *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, 2018, pp. 2482-2491.
- [20] M. Li, J. Liu, W. Yang, X. Sun, and Z. Guo, "Structure-revealing low-light image enhancement via robust retinex model," *IEEE Trans. on Image Processing*, vol. 27, no. 6, pp. 2828-2841, June 2018.
- [21] X. Ren, M. Li, W.-H. Cheng, and J. Liu, "Joint enhancement and denoising method via sequential decomposition," in *IEEE Int'l Symposium on Circuits and Systems*, May 2018, pp. 1-5.
- [22] J. Yang, X. Jiang, C. Pan, and C.-L. Liu, "Enhancement of low light level images with coupled dictionary learning," in *Proc. IEEE Int'l Conf. Pattern Recognition*, Dec 2016, pp. 751-756.
- [23] K. G. Lore, A. Akintayo, and S. Sarkar, "Llnet: A deep autoencoder approach to natural low-light image enhancement," *Pattern Recognition*, vol. 61, pp. 650-662, 2017.
- [24] L. Shen, Z. Yue, F. Feng, Q. Chen, S. Liu, and J. Ma, "MSR-net: Low-light Image Enhancement Using Deep Convolutional Network," *ArXiv e-prints*, arXiv:1711.02488, November 2017.
- [25] Yifan Jiang, Xinyu Gong, Ding Liu, Yu Cheng, Chen Fang, Xiaohui Shen, Jianchao Yang, Pan Zhou, Zhangyang Wang, "EnlightenGAN: Deep Light Enhancement without Paired Supervision," *ArXiv e-prints*, arXiv:1906.06972, June 2019.
- [26] Yueyu Hu, Wenhan Yang, Jiaying Liu, "Coarse-to-Fine Hyper-Prior Modeling for Learned Image Compression," in *Proc. AAAI Conference on Artificial Intelligence*, Feb. 2020.
- [27] Jiaying Liu, Sifeng Xia, and Wenhan Yang, "Deep Reference Generation with Multi-Domain Hierarchical Constraints for Inter Prediction," in *IEEE Transactions on Multimedia*, Dec. 2019.
- [28] Jiaying Liu, Sifeng Xia, Wenhan Yang, Mading Li, and Dong Liu, "One-for-All: Grouped Variation Network Based Fractional Interpolation in Video Coding," in *IEEE Trans. on Image Processing*, Vol.28, No.5, pp.2140-2151, May 2019.
- [29] Xiaoshuai Zhang, Wenhan Yang, Yueyu Hu, Jiaying Liu, "DMCNN: Dual-Domain Multi-Scale Convolutional Neural Network for Compression Artifacts Removal," in *Proc. IEEE International Conference on Image Processing*, Oct. 2018.
- [30] Dezha Wang, Sifeng Xia, Wenhan Yang, Yueyu Hu, Jiaying Liu, "Partition Tree Guided Progressive Rethinking Network for in-Loop Filtering of HEVC," in *Proc. IEEE International Conference on Image Processing*, Sep. 2019.
- [31] B. Wen, Y. Zhu, R. Subramanian, T. Ng, X. Shen and S. Winkler, "COVERAGE - A novel database for copy-move forgery detection," in *Proc. IEEE Int'l Conf. on Image Processing*, 2016, pp. 161-165.
- [32] K. He, X. Zhang, S. Ren, and J. Sun, "Deep residual learning for image recognition," in *Proc. IEEE Int'l Conf. Computer Vision and Pattern Recognition*, June 2016, pp. 770-778.
- [33] G. Papandreou, L. Chen, K. P. Murphy, and A. L. Yuille, "Weakly-and semi-supervised learning of a deep convolutional network for semantic image segmentation," in *Proc. IEEE Int'l Conf. Computer Vision*, Dec 2015, pp. 1742-1750.
- [34] K. Simonyan and A. Zisserman, "Two-stream convolutional networks for action recognition in videos," in *Proc. Int'l Conf. on Neural Information Processing Systems*, 2014, pp. 568-576.
- [35] R. Alp Gler, N. Neverova, and I. Kokkinos, "Densepose: Dense human pose estimation in the wild," in *Proc. IEEE Int'l Conf. Computer Vision and Pattern Recognition*, June 2018, pp. 7297-7306.
- [36] D. Bradbury, "How autonomous vehicles will navigate bad weather remains foggy," <https://www.forbes.com/sites/centurylink/2016/11/29/how-autonomous-vehicles-will-navigate-bad-weather-remains-foggy/#2c9b3d248662>, forbes November 29, 2016.
- [37] A. Ashe, "Ticketbuster: Snow causes red-light camera to issue ticket in error," <https://wtop.com/news/2014/04/ticketbuster-snow-causes-red-light-camera-to-issue-ticket-in-error/>, wtop April 18, 2014.
- [38] Y. Liu, G. Zhao, B. Gong, Y. Li, R. Raj, N. Goel, S. Kesav, S. Gottimukkala, Z. Wang, W. Ren *et al.*, "Improved techniques for learning to dehaze and beyond: A collective study," *arXiv preprint, arXiv:1807.00202*, 2018.
- [39] H. Zhang, J. Yang, Y. Zhang, N. M. Nasrabadi, and T. S. Huang, "Close the loop: Joint blind image restoration and recognition with sparse representation prior," in *Proc. IEEE Int'l Conf. Computer Vision*, 2011, pp. 770-777.
- [40] D. Liu, B. Cheng, Z. Wang, H. Zhang, and T. S. Huang, "Enhance visual recognition under adverse conditions via deep networks," *arXiv preprint, arXiv:1712.07732*, 2017.
- [41] D. Liu, B. Wen, J. Jiao, X. Liu, Z. Wang, and T. S. Huang, "Connecting image denoising and high-level vision tasks via deep learning," *arXiv preprint, arXiv:1809.01826*, 2018.
- [42] D. Liu, B. Wen, X. Liu, Z. Wang, and T. S. Huang, "When image denoising meets high-level vision tasks: A deep learning approach," *arXiv preprint, arXiv:1706.04284*, 2017.
- [43] B. Cheng, Z. Wang, Z. Zhang, Z. Li, D. Liu, J. Yang, S. Huang, and T. S. Huang, "Robust emotion recognition from low quality and low bit rate video: A deep learning approach," in *Proc. Int'l Conf. on Affective Computing and Intelligent Interaction*, 2017, pp. 65-70.
- [44] L. Stasiak, A. Pacut, and R. Vicente-Garcia, "Face tracking and recognition in low quality video sequences with the use of particle filtering," in *proc. Annual*

- Int'l Carnahan Conf. on Security Technology*, 2009, pp. 126-133.
- [45] V. Vašek, V. Franc, and M. Urban, "License plate recognition and super-resolution from low-resolution videos by convolutional neural networks," in *proc. British Machine Vision Conference*, 2018.
 - [46] Y. li Tian, "Evaluation of face resolution for expression analysis," in *proc. IEEE Int'l Conf. Computer Vision and Pattern Recognition Workshop*, 2004, pp. 82-82.
 - [47] C. Shan, S. Gong, and P. McOwan, "Recognizing facial expressions at low resolution," in *proc. IEEE Conf. on Advanced Video and Signal Based Surveillance*, 2005.
 - [48] R. Prabhu, X. Yu, Z. Wang, D. Liu, and A. Jiang, "U-finger: Multi-scale dilated convolutional network for fingerprint image denoising and inpainting," *arXiv preprint arXiv:1807.10993*, 2018.
 - [49] M. Bevilacqua, A. Roumy, C. Guillemot, and M. line Alberi Morel, "Low-complexity single-image super-resolution based on nonnegative neighbor embedding," in *proc. the British Machine Vision Conf.*, 2012, pp. 135.1-135.10.
 - [50] R. Zeyde, M. Elad, and M. Protter, "On single image scale-up using sparse-representations," in *proc. the Int'l Conf. on Curves and Surfaces*, 2012, pp. 711-730.
 - [51] H. R. Sheikh, M. F. Sabir, and A. C. Bovik, "A statistical evaluation of recent full reference image quality assessment algorithms," *IEEE Trans. on Image Processing*, vol. 15, no. 11, pp. 3440-3451, Nov 2006.
 - [52] D. Martin, C. Fowlkes, D. Tal, and J. Malik, "A database of human segmented natural images and its application to evaluating segmentation algorithms and measuring ecological statistics," in *Proc. IEEE Int'l Conf. Computer Vision*, vol. 2, July 2001, pp. 416-423.
 - [53] J. Huang, A. Singh, and N. Ahuja, "Single image super-resolution from transformed self-exemplars," in *Proc. IEEE Int'l Conf. Computer Vision and Pattern Recognition*, 2015, pp. 5197-5206.
 - [54] R. Timofte, E. Agustsson, L. Van Gool, M.-H. Yang, and L. Zhang, "NTIRE 2017 challenge on single image super-resolution: Methods and results," in *Proc. IEEE Int'l Conf. Computer Vision and Pattern Recognition Workshop*, 2017, pp. 1110-1121.
 - [55] A. Fujimoto, T. Ogawa, K. Yamamoto, Y. Matsui, T. Yamasaki, and K. Aizawa, "Manga109 dataset and creation of metadata," in *proc. Int'l Workshop on coMics Analysis, Processing and Understanding*, 2016, pp. 1-5.
 - [56] J. Xu, H. Li, Z. Liang, D. Zhang, and L. Zhang, "Real-world Noisy Image Denoising: A New Benchmark," *arXiv e-prints, arXiv:1804.02603*, Apr 2018.
 - [57] Tobias Plotz and Stefan Roth, "Benchmarking Denoising Algorithms With Real Photographs," in *Proc. IEEE Int'l Conf. Computer Vision and Pattern Recognition*, July 2017.
 - [58] C. Chen, Q. Chen, J. Xu, and V. Koltun, "Learning to See in the Dark," in *Proc. IEEE Int'l Conf. Computer Vision and Pattern Recognition*, pp. 3291-3300.
 - [59] C. Wei, W. Wang, W. Yang, and J. Liu, "Deep Retinex decomposition for low-light enhancement," in *Proc. British Machine Vision Conference*, 2018, p. 155.
 - [60] Y. Zhang, L. Ding, and G. Sharma, "Hazerd: an outdoor scene dataset and benchmark for single image dehazing," in *Proc. IEEE Int'l Conf. Image Processing*, 2017, pp. 3205-3209.
 - [61] C. O. Ancuti, C. Ancuti, R. Timofte, and C. De Vleeschouwer, "O-HAZE: a dehazing benchmark with real hazy and haze-free outdoor images," *arXiv e-prints, arXiv:1804.05101*, April 2018.
 - [62] —, "I-HAZE: a dehazing benchmark with real hazy and haze-free indoor images," *arXiv e-prints, arXiv:1804.05091*, April 2018.
 - [63] H. Zhang, V. Sindagi, and V. M. Patel, "Image de-raining using a conditional generative adversarial network," *arXiv e-prints, arXiv:1701.05957*, 2017.
 - [64] R. Qian, R. T. Tan, W. Yang, J. Su, and J. Liu, "Attentive generative adversarial network for raindrop removal from a single image," in *Proc. IEEE Conf. on Computer Vision and Pattern Recognition*, 2018, pp. 2482-2491.
 - [65] M. Grgic, K. Delac, and S. Grgic, "SCface — surveillance cameras face database," *Multimedia Tools Appl.*, vol. 51, no. 3, pp. 863-879, February 2011.
 - [66] J. Shao, C. C. Loy, and X. Wang, "Scene-Independent Group Profiling in Crowd," in *Proc. IEEE Int'l Conf. Computer Vision and Pattern Recognition*, 2014, pp. 2227-2234.
 - [67] X. Zhu, C. C. Loy, and S. Gong, "Video synopsis by heterogeneous multi-source correlation," in *Proc. IEEE Int'l Conf. Computer Vision*, Dec 2013, pp. 81-88.
 - [68] S. Oh, A. Hoogs, A. Perera, N. Cuntoor, C. Chen, J. T. Lee, S. Mukherjee, J. K. Aggarwal, H. Lee, L. Davis, E. Swears, X. Wang, Q. Ji, K. Reddy, M. Shah, C. Vondrick, H. Pirsiavash, D. Ramanan, J. Yuen, A. Torralba, B. Song, A. Fong, A. Roy-Chowdhury, and M. Desai, "A large-scale benchmark dataset for event recognition in surveillance video," in *Proc. IEEE Int'l Conf. Computer Vision and Pattern Recognition*, June 2011, pp. 3153-3160.
 - [69] M. Mueller, N. Smith, and B. Ghanem, "A benchmark and simulator for UAV tracking," in *Proc. European Conf. Computer Vision*, 2016, pp. 445-461.
 - [70] B. Z. Yao, X. Yang, and S.-C. Zhu, "Introduction to a large-scale general purpose ground truth database: Methodology, annotation tool and benchmarks," in *Proc. Energy Minimization Methods in Computer Vision and Pattern Recognition*, 2007, pp. 169-183.
 - [71] H. Nada, V. A. Sindagi, H. Zhang, and V. M. Patel, "Pushing the Limits of Unconstrained Face Detection: a Challenge Dataset and Baseline Results," *arXiv e-prints, arXiv:1804.10275*, Apr. 2018.
 - [72] B. Li, W. Ren, D. Fu, D. Tao, D. Feng, W. Zeng, and Z. Wang, "Benchmarking single-image dehazing and beyond," *IEEE Trans. on Image Processing*, vol. 28, no. 1, pp. 492-505, 2019.
 - [73] Y. P. Loh and C. S. Chan, "Getting to know low-light images with the exclusively dark dataset," *Computer Vision and Image Understanding*, vol. 178, pp. 30-42, 2019.
 - [74] S. Li, I. B. Araujo, W. Ren, Z. Wang, E. K. Tokuda, R. H. Junior, R. Cesar-Junior, J. Zhang, X. Guo, and X. Cao, "Single image deraining: A comprehensive benchmark analysis," in *Proc. Conf. on Computer Vision and Pattern Recognition (CVPR)*, 2019, pp. 3833-3842.
 - [75] Z. Wang, H. Li, Q. Ling, and W. Li, "Robust temporal-spatial decomposition and its applications in video processing," *IEEE Trans. on Circuits and Systems for Video Technology*, vol. 23, no. 3, pp. 387-400, 2013.
 - [76] H. Li, Z. Lu, Z. Wang, Q. Ling, and W. Li, "Detection of blotch and scratch in video based on video decomposition," *IEEE Trans. on Circuits and Systems for Video Technology*, vol. 23, no. 11, pp. 1887-1900, 2013.
 - [77] J. Ren, J. Liu, and Z. Guo, "Context-aware sparse decomposition for image denoising and super-resolution," *IEEE Trans. on Image Processing*, vol. 22, no. 4, pp. 1456-1469, 2013.
 - [78] Z. Wang, D. Liu, S. Chang, Q. Ling, Y. Yang, and T. S. Huang, "D3: Deep dual-domain based fast restoration of JPEG-compressed images," in *Proc. IEEE Int'l Conf. Computer Vision and Pattern Recognition*, 2016, pp. 2764-2772.
 - [79] J. Liu, S. Yang, Y. Fang, and Z. Guo, "Structure-guided image inpainting using homography transformation," *IEEE Trans. on Multimedia*, vol. 20, no. 12, pp. 3252-3265, 2018.
 - [80] Y. Yan, W. Ren, Y. Guo, R. Wang, and X. Cao, "Image deblurring via extreme channels prior," in *Proc. IEEE Int'l Conf. Computer Vision and Pattern Recognition*, 2017, pp. 4003-4011.
 - [81] W. Ren, J. Pan, X. Cao, and M.-H. Yang, "Video deblurring via semantic segmentation and pixel-wise non-linear kernel," in *Proc. IEEE Int'l Conf. on Computer Vision*, 2017, pp. 1077-1085.
 - [82] Z. Wang, Y. Yang, Z. Wang, S. Chang, J. Yang, and T. S. Huang, "Learning super-resolution jointly from external and internal examples," *IEEE Trans. on Image Processing*, vol. 24, no. 11, pp. 4359-4371, 2015.
 - [83] Z. Wang, Y. Yang, Z. Wang, S. Chang, W. Han, J. Yang, and T. Huang, "Self-tuned deep super resolution," in *Proc. IEEE Conf. on Computer Vision and Pattern Recognition Workshops*, 2015, pp. 1-8.
 - [84] D. Liu, Z. Wang, Y. Fan, X. Liu, Z. Wang, S. Chang, and T. Huang, "Robust video super-resolution with learned temporal dynamics," in *Proc. IEEE Int'l Conf. Computer Vision*, 2017, pp. 2526-2534.
 - [85] D. Liu, Z. Wang, Y. Fan, X. Liu, Z. Wang, S. Chang, X. Wang, and T. S. Huang, "Learning temporal dynamics for video super-resolution: A deep learning approach," *IEEE Trans. on Image Processing*, vol. 27, no. 7, pp. 3432-3445, July 2018.
 - [86] Z. Yu, H. Li, Z. Wang, Z. Hu, and C. W. Chen, "Multi-level video frame interpolation: Exploiting the interaction among different levels," *IEEE Trans. on Circuits and Systems for Video Technology*, vol. 23, no. 7, pp. 1235-1248, 2013.
 - [87] R. Fattal, "Single image dehazing," *ACM Trans. Graph.*, vol. 27, no. 3, pp. 72:1-72:9, August 2008.
 - [88] Q. Zhu, J. Mai, and L. Shao, "A fast single image haze removal algorithm using color attenuation prior," *IEEE Trans. on Image Processing*, vol. 24, no. 11, pp. 3522-3533, Nov 2015.
 - [89] D. Berman, T. Treibitz, and S. Avidan, "Non-local Image Dehazing," in *Proc. IEEE Int'l Conf. Computer Vision and Pattern Recognition*, 2016, pp. 1674-1682.
 - [90] J. Zhou and F. Zhou, "Single image dehazing motivated by retinex theory," in *Proc. Int'l Symposium on Instrumentation and Measurement, Sensor Network and Automation*, Dec 2013, pp. 243-247.
 - [91] D. Nair, P. A. Kumar, and P. Sankaran, "An effective surround filter for image dehazing," in *Proc. Int'l Conf. on Interdisciplinary Advances in Applied Computing*, 2014, pp. 20:1-20:6.
 - [92] L. Kratz and K. Nishino, "Factorizing scene albedo and depth from a single foggy image," in *Proc. IEEE Int'l Conf. Computer Vision*, Sep. 2009, pp. 1701-1708.
 - [93] K. Nishino, L. Kratz, and S. Lombardi, "Bayesian defogging," *Int'l Journal of Computer Vision*, vol. 98, no. 3, pp. 263-278, July 2012.
 - [94] Y. Li, R. T. Tan, and M. S. Brown, "Nighttime haze removal with glow and multiple light colors," in *Proc. IEEE Int'l Conf. Computer Vision*, Dec 2015, pp. 226-234.
 - [95] J. Zhang, Y. Cao, S. Fang, Y. Kang, and C. W. Chen, "Fast Haze Removal for Nighttime Image Using Maximum Reflectance Prior," in *Proc. IEEE Int'l Conf. Computer Vision and Pattern Recognition*, pp. 7016-7024.
 - [96] B. Li, X. Peng, Z. Wang, J. Xu, and D. Feng, "An all-in-one network for dehazing and beyond," *arXiv preprint, arXiv:1707.06543*, 2017.
 - [97] —, "End-to-end united video dehazing and detection," in *Proc. AAAI Conference on Artificial Intelligence*, Feb. 2018.
 - [98] W. Ren, J. Zhang, X. Xu, L. Ma, X. Cao, G. Meng, and W. Liu, "Deep video dehazing with semantic segmentation," *IEEE Trans. on Image Processing*, vol. 28, no. 4, pp. 1895-1908, April 2019.
 - [99] S. M. Pizer, R. E. Johnston, J. P. Erickson, B. C. Yankaskas, and K. E. Muller, "Contrast-limited adaptive histogram equalization: speed and effectiveness," in *Proc. Conf. on Visualization in Biomedical Computing*, May 1990, pp. 337-345.
 - [100] M. Abdullah-Al-Wadud, M. H. Kabir, M. A. A. Dewan, and O. Chae, "A dynamic histogram equalization for image contrast enhancement," *IEEE Transactions on Consumer Electronics*, vol. 53, no. 2, pp. 593-600, May 2007.
 - [101] L. Li, R. Wang, W. Wang, and W. Gao, "A low-light image enhancement method for both denoising and contrast enlarging," in *Proc. IEEE Int'l Conf. Image Processing*, Sept 2015, pp. 3730-3734.
 - [102] X. Zhang, P. Shen, L. Luo, L. Zhang, and J. Song, "Enhancement and noise reduction of very low light level images," in *Proc. IEEE Int'l Conf. Pattern Recognition*, Nov 2012, pp. 2034-2037.
 - [103] E. H. Land, "The retinex theory of color vision," *Sci. Amer.*, pp. 108-128, 1977.
 - [104] D. J. Jobson, Z. Rahman, and G. A. Woodell, "Properties and performance of a center/surround retinex," *IEEE Trans. on Image Processing*, vol. 6, no. 3, pp.

- 451–462, Mar 1997.
- [105] —, “A multiscale retinex for bridging the gap between color images and the human observation of scenes,” *IEEE Trans. on Image Processing*, vol. 6, no. 7, pp. 965–976, July 1997.
- [106] S. Wang, J. Zheng, H. M. Hu, and B. Li, “Naturalness preserved enhancement algorithm for non-uniform illumination images,” *IEEE Trans. on Image Processing*, vol. 22, no. 9, pp. 3538–3548, Sept 2013.
- [107] X. Fu, D. Zeng, Y. Huang, Y. Liao, X. Ding, and J. Paisley, “A fusion-based enhancing method for weakly illuminated images,” *Signal Processing*, vol. 129, pp. 82–96, 2016.
- [108] X. Guo, Y. Li, and H. Ling, “Lime: Low-light image enhancement via illumination map estimation,” *IEEE Trans. on Image Processing*, vol. 26, no. 2, pp. 982–993, Feb 2017.
- [109] X. Fu, D. Zeng, Y. Huang, X. P. Zhang, and X. Ding, in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2016, pp. 2782–2790.
- [110] J. Cai, S. Gu, and L. Zhang, “Learning a deep single image contrast enhancer from multi-exposure images,” *IEEE Trans. on Image Processing*, April 2018.
- [111] Y. Jiang, X. Gong, D. Liu, Y. Cheng, C. Fang, X. Shen, J. Yang, P. Zhou, and Z. Wang, “Enlighten: Deep light enhancement without paired supervision,” *arXiv preprint, arXiv:1906.06972*, 2019.
- [112] L. W. Kang, C. W. Lin, and Y. H. Fu, “Automatic single-image-based rain streaks removal via image decomposition,” *IEEE Trans. on Image Processing*, vol. 21, no. 4, pp. 1742–1755, April 2012.
- [113] J. H. Kim, C. Lee, J. Y. Sim, and C. S. Kim, “Single-image deraining using an adaptive nonlocal means filter,” in *Proc. IEEE Int’l Conf. Image Processing*, Sept 2013, pp. 914–917.
- [114] K. Garg and S. K. Nayar, “Photorealistic rendering of rain streaks,” in *ACM Trans. Graphics*, vol. 25, no. 3, 2006, pp. 996–1002.
- [115] —, “Detection and removal of rain from videos,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2004, pp. 1–1.
- [116] —, “Vision and rain,” *Int’l Journal of Computer Vision*, vol. 75, no. 1, pp. 3–27, 2007.
- [117] —, “When does a camera see rain?” in *Proc. IEEE Int’l Conf. Computer Vision*, vol. 2, 2005, pp. 1067–1074.
- [118] X. Zhang, H. Li, Y. Qi, W. K. Leow, and T. K. Ng, “Rain removal in video by combining temporal and chromatic properties,” in *Proc. IEEE Int’l Conf. Multimedia and Expo*, 2006, pp. 461–464.
- [119] P. Liu, J. Xu, J. Liu, and X. Tang, “Pixel based temporal analysis using chromatic property for removing rain from videos,” in *Computer and Information Science*, 2009.
- [120] P. C. Barnum, S. Narasimhan, and T. Kanade, “Analysis of rain and snow in frequency space,” *International Journal of Computer Vision*, vol. 86, no. 2–3, pp. 256–274, 2010.
- [121] V. Santhaseelan and V. K. Asari, “Utilizing local phase information to remove rain from video,” *Int’l Journal of Computer Vision*, vol. 112, no. 1, pp. 71–89, March 2015.
- [122] N. Brewer and N. Liu, “Using the shape characteristics of rain to identify and remove rain from video,” in *Joint IAPR Int’l Workshops on SPR and SSPR*, 2008, pp. 451–458.
- [123] J. Bossu, N. Hautière, and J.-P. Tarel, “Rain or snow detection in image sequences through use of a histogram of orientation of streaks,” *Int’l Journal of Computer Vision*, vol. 93, no. 3, pp. 348–367, 2011.
- [124] T.-X. Jiang, T.-Z. Huang, X.-L. Zhao, L.-J. Deng, and Y. Wang, “A novel tensor-based video rain streaks removal approach via utilizing discriminatively intrinsic priors,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2017, pp. 2818–2827.
- [125] J. Chen and L. P. Chau, “A rain pixel recovery algorithm for videos with highly dynamic scenes,” *IEEE Trans. on Image Processing*, vol. 23, no. 3, pp. 1097–1104, March 2014.
- [126] A. K. Tripathi and S. Mukhopadhyay, “A probabilistic approach for detection and removal of rain from videos,” *IETE Journal of Research*, vol. 57, no. 1, pp. 82–91, 2011.
- [127] —, “Video post processing: low-latency spatiotemporal approach for detection and removal of rain,” *IET Image Processing*, vol. 6, no. 2, pp. 181–196, March 2012.
- [128] W. Ren, J. Tian, Z. Han, A. Chan, and Y. Tang, “Video desnowing and deraining based on matrix decomposition,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2017, pp. 2838–2847.
- [129] M. Li, Q. Xie, Q. Zhao, W. Wei, S. Gu, J. Tao, and D. Meng, “Video rain streak removal by multiscale convolutional sparse coding,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2018, pp. 6644–6653.
- [130] W. Wei, L. Yi, Q. Xie, Q. Zhao, D. Meng, and Z. Xu, “Should we encode rain streaks in video as deterministic or stochastic?” in *Proc. IEEE Int’l Conf. Computer Vision*, Oct 2017, pp. 2535–2544.
- [131] J. H. Kim, J. Y. Sim, and C. S. Kim, “Video deraining and desnowing using temporal correlation and low-rank matrix completion,” *IEEE Trans. on Image Processing*, vol. 24, no. 9, pp. 2658–2670, Sept 2015.
- [132] J. Chen, C.-H. Tan, J. Hou, L.-P. Chau, and H. Li, “Robust video content alignment and compensation for rain removal in a CNN framework,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2018, pp. 6286–6295.
- [133] J. Liu, W. Yang, S. Yang, and Z. Guo, “Erase or fill? deep joint recurrent rain removal and reconstruction in videos,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2018, pp. 3233–3242.
- [134] A. Torralba, R. Fergus, and W. T. Freeman, “80 million tiny images: A large data set for nonparametric object and scene recognition,” *IEEE Trans. on Pattern Analysis and Machine Intelligence*, vol. 30, no. 11, pp. 1958–1970, Nov. 2008.
- [135] W. W. Zou and P. C. Yuen, “Very low resolution face recognition problem,” *IEEE Trans. on Image Processing*, vol. 21, no. 1, pp. 327–340, Jan. 2012.
- [136] S. Karahan, M. Kilinc Yildirim, K. Kirtac, F. S. Rende, G. Butun, and H. K. Ekenel, “How image degradations affect deep CNN-based face recognition?” in *Int’l Conf. of the Biometrics Special Interest Group*, Sep. 2016.
- [137] A. Dutta, R. Veldhuis, and L. Spreeuwiers, “The impact of image quality on the performance of face recognition,” in *Symposium on Information Theory in the Benelux and Joint WIC/IEEE Symposium on Information Theory and Signal Processing in the Benelux*, 2012.
- [138] R. Fergus, B. Singh, A. Hertzmann, S. T. Roweis, and W. T. Freeman, “Removing camera shake from a single photograph,” in *ACM Trans. on Graphics*, 2006, pp. 787–794.
- [139] J. Yang, J. Wright, T. S. Huang, and Y. Ma, “Image Super-Resolution Via Sparse Representation,” in *IEEE Trans. on Image Processing*, vol. 19, no. 11, pp. 2861–2873, Nov. 2010.
- [140] C. Dong, Y. Deng, C. C. Loy, and X. Tang, “Compression Artifacts Reduction by a Deep Convolutional Network,” in *Proc. IEEE Int’l Conf. on Computer Vision*, 2015, pp. 576–584.
- [141] P. H. Hennings-Yeomans, S. Baker, and B. V. K. V. Kumar, “Simultaneous super-resolution and feature extraction for recognition of low-resolution faces,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2008, pp. 1–8.
- [142] Z. Wang, S. Chang, Y. Yang, D. Liu, and T. S. Huang, “Studying very low resolution recognition using deep networks,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2016, pp. 4792–4800.
- [143] R. G. VidalMata, S. Banerjee, B. RichardWebster, M. Albright, P. Davalos, S. McCloskey, B. Miller, A. Tambo, S. Ghosh, S. Nagesh et al., “Bridging the gap between computational photography and visual recognition,” *arXiv preprint, arXiv:1901.09482*, 2019.
- [144] K. He, G. Gkioxari, P. Dollár, and R. Girshick, “Mask R-CNN,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2017, pp. 2980–2988.
- [145] Cheng Chi, Shifeng Zhang, Junliang Xing, Zhen Lei, Stan Z. Li, Xudong Zou, “Selective refinement network for high performance face detection,” in *Proc. AAAI Conference on Artificial Intelligence*, Feb. 2019.
- [146] T.-Y. Lin, P. Goyal, R. Girshick, K. He, and P. Dollár, “Focal loss for dense object detection,” *IEEE Trans. on Pattern Analysis and Machine Intelligence*, 2017, pp. 2999–3007.
- [147] J. Redmon and A. Farhadi, “Yolov3: An incremental improvement,” *arXiv e-prints, arXiv:1804.02767*, 2018.
- [148] T.-Y. Lin, P. Dollár, R. Girshick, K. He, B. Hariharan, and S. Belongie, “Feature pyramid networks for object detection,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2017, pp. 936–944.
- [149] H. Zhang and V. M. Patel, “Densely connected pyramid dehazing network,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2018, pp. 3194–3203.
- [150] J. Li, Y. Wang, C. Wang, Y. Tai, J. Qian, J. Yang, C. Wang, J. Li, and F. Huang, “DSFD: Dual shot face detector,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2019, pp. 5055–5064.
- [151] X. Tang, D. K. Du, Z. He, and J. Liu, “Pyramidbox: A context-assisted single shot face detector,” in *Proc. European Conf. on Computer Vision*, 2018, pp. 812–828.
- [152] M. Najibi, P. Samangouei, R. Chellappa, and L. S. Davis, “Ssh: Single stage headless face detector,” in *Proc. IEEE Int’l Conf. Computer Vision*, 2017.
- [153] H. Jiang and E. G. Learned-Miller, “Face detection with the faster R-CNN,” *IEEE Int’l Conf. on Automatic Face and Gesture Recognition*, 2017.
- [154] Z. Ying, G. Li, and W. Gao, “A Bio-Inspired Multi-Exposure Fusion Framework for Low-light Image Enhancement,” *arXiv e-prints, arXiv:1711.00591*, November 2017.
- [155] X. Dong, G. Wang, Y. Pang, W. Li, J. Wen, W. Meng, and Y. Lu, “Fast efficient algorithm for enhancement of low lighting video,” in *Proc. IEEE Int’l Conf. Multimedia and Expo*, 2011, pp. 1–6.
- [156] S. Yang, P. Luo, C. C. Loy, and X. Tang, “WIDER FACE: A Face Detection Benchmark,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2016, pp. 5525–5533.
- [157] S. Ren, K. He, R. Girshick, and J. Sun, “Faster r-cnn: Towards real-time object detection with region proposal networks,” in *Advances in Neural Information Processing Systems*, 2015, pp. 91–99.
- [158] W. Liu, D. Anguelov, D. Erhan, C. Szegedy, S. Reed, C.-Y. Fu, and A. C. Berg, “Ssd: Single shot multibox detector,” in *European conf. on computer vision*, 2016, pp. 21–37.
- [159] S. Xie, R. Girshick, P. Dollár, Z. Tu, and K. He, “Aggregated residual transformations for deep neural networks,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2017, pp. 1492–1500.
- [160] P. Izmailov, D. Podoprikin, T. Garipov, D. Vetrov, and A. G. Wilson, “Averaging weights leads to wider optima and better generalization,” *arXiv preprint, arXiv:1803.05407*, 2018.
- [161] Z. Cai and N. Vasconcelos, “Cascade R-CNN: Delving into high quality object detection,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2018, pp. 6154–6162.
- [162] O. Kupyn, V. Budzan, M. Mykhailych, D. Mishkin, and J. Matas, “Deblurgan: Blind motion deblurring using conditional adversarial networks,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2018, pp. 8183–8192.
- [163] F. Drago, K. Myszkowski, T. Annen, and N. Chiba, “Adaptive logarithmic mapping for displaying high contrast scenes,” in *Computer Graphics Forum*, 2003.
- [164] O. Ronneberger, P. Fischer, and T. Brox, “U-Net: Convolutional networks for biomedical image segmentation,” in *Proc. Int’l Conf. on Medical image computing and computer-assisted intervention*, 2015, pp. 234–241.
- [165] K. Zhang, W. Zuo, and L. Zhang, “Deep plug-and-play super-resolution for

arbitrary blur kernels,” in *Proc. IEEE Int’l Conf. Computer Vision and Pattern Recognition*, 2019, pp. 1671–1681.