

EVIDENCE BASED VERIFICATION OF CAMERA LENS CLEANING OUTCOMES IN INTELLIGENT VISION SYSTEMS

Abdulazizxon Axmedov Ganijon O'g'li

PhD student, Namangan State University, Namangan, Uzbekistan

Email: abdulazizahmedov97@gmail.com

Abstract:

An automated camera cleaner can complete its commanded action while contamination still degrades the image. This paper proposes a procedure for verifying visual recovery after a cleaning intervention. The procedure associates observations with a specific intervention, excludes images captured before mechanical and optical settling, preserves comparable pre-cleaning evidence, and requires a bounded sequence of acceptable observations. Residual contamination, obstruction of designated critical regions, and an independently validated image-usability check jointly determine acceptance. An exact calculation with an assumed binary Markov error model examines the effect of temporally correlated false-pass decisions. For a marginal false-pass probability of 0.10, three consecutive passes within 30 observations give a false-clear probability of 2.506% under independence and 38.149% when lag-one correlation is 0.70. These are analytical scenario results, not measurements from a camera or cleaning device. The analysis shows why persistence alone cannot establish recovery and why the observation horizon must be specified. The contribution is an intervention-level verification procedure and its conditional reliability analysis; physical cleaning efficiency and real-world perception recovery remain to be experimentally evaluated.

Keywords: camera lens contamination, cleaning verification, visual recovery, temporal dependence, camera health monitoring, finite observation horizon.

Introduction

Camera cleaning has two observable outcomes: the mechanism completes an action, and the resulting image becomes usable for its intended task. These outcomes can differ. A washer may discharge insufficient fluid, a wiper may redistribute contamination, or a cleaning cycle may leave a transparent film that affects contrast. A control system therefore needs evidence from the image after the intervention before it updates the camera's recovery status.

Experimental work by Kim et al. demonstrates that droplet contamination affects both optical image quality and object detection, with effects depending on droplet characteristics and surface properties [1]. This motivates checking more than the apparent disappearance of a segmented region. ISO 24650:2024 addresses assessment of cleaning-system efficiency through removal of contamination from the sensor's front surface; its published scope explicitly excludes sensor-performance indicators [2]. An image-based recovery decision consequently addresses a different assessment question from surface-cleaning efficiency.

The author's published camera-monitoring framework combines image quality assessment with lens-defect analysis and discusses response mechanisms [3]. The present work develops one specific subsequent decision: what evidence permits the system to record an acceptable visual condition after an identified cleaning action? It does not introduce a new segmentation network or claim improved defect-localization accuracy. Its contributions are a precise verification procedure, worked boundary cases, and an exact finite-horizon analysis showing how correlated errors weaken a consecutive-observation acceptance rule.

Related Work and Study Scope

TiledSoilingNet estimates soiling coverage at the tile level in automotive surround-view images [4]. Such outputs can support residual-contamination assessment, but a coverage estimate alone does not establish when the image was captured relative to a cleaning cycle or whether a usable condition persisted afterward. The proposed verifier accepts detector outputs through a defined interface and treats their correctness as an assumption to be evaluated separately.

The Desoiling Dataset provides paired clean and soiled video information and investigates image restoration [5]. Paired observations are useful for studying visual degradation, but validating an intervention verifier additionally requires action timestamps, actuator outcomes, settling periods, and independent recovery labels. Image restoration also changes image content computationally; it cannot by itself demonstrate that a physical obstruction was removed.

Confidence must be interpreted carefully. Guo et al. show that neural-network confidence can be poorly calibrated [6]. Even calibrated class confidence describes the reliability of a specified prediction under its calibration conditions. It does not measure lens transparency, spatial resolution, or task performance directly. Accordingly, the procedure keeps defect estimates and image-usability evidence separate. Their errors may nevertheless share causes, so the analysis does not assume that combining them makes the resulting decisions independent.

ISO 21448:2022 provides a broader framework for addressing functional insufficiencies in systems that depend on situational awareness [7]. The present procedure is a proposed monitoring component, not a demonstration of compliance with that standard or an authorization to resume an automated-driving function. Vehicle-level operating decisions remain outside the study. No physical cleaner, trained recovery classifier, or driving dataset was evaluated for this manuscript.

Methodology

This study employed a methodological and analytical research design to develop and evaluate an evidence-based procedure for verifying camera-lens cleaning outcomes in intelligent vision systems. Rather than measuring the mechanical efficiency of a physical cleaning device, the study focused on determining whether post-intervention visual evidence is sufficient to classify a camera condition as

acceptably recovered. The proposed methodology integrates intervention tracking, observation eligibility criteria, residual-contamination assessment, image-usability verification, bounded consecutive confirmation, and probabilistic analysis of false-clear events [8].

Each cleaning attempt was treated as an independent intervention episode and assigned a unique identifier. Before the cleaning command was initiated, the system retained a checkpoint containing the most recent valid contamination measurements, acquisition timestamp, camera configuration, and predefined acceptance thresholds. Following completion of the cleaning action, observations were excluded during a configurable settling interval to prevent temporary disturbances caused by spray, wiper movement, exposure changes, or other transient optical effects from influencing the verification decision.

Two contamination indicators were considered. The first, (S), represents the estimated proportion of the monitored image affected by contamination, while (C) represents the maximum contaminated proportion within predefined task-critical regions. These indicators were complemented by a Boolean image-usability variable, (U), representing independently validated evidence that the resulting image remained suitable for its intended visual task. An observation was considered eligible only when it was acquired after the settling period, belonged to the active intervention, had not previously been counted, matched the required camera configuration, and contained all measurements necessary for verification [9].

Recovery was verified when (N) consecutive eligible observations satisfied all predefined acceptance conditions within a bounded verification horizon. A valid failing observation or missing required evidence interrupted the consecutive sequence. The verification episode was additionally constrained by a wall-clock deadline and a maximum number (H) of eligible observations. Episodes that failed to satisfy the complete acceptance criteria were classified, where sufficient comparable evidence existed, as partial improvement, unchanged, worsened, or inconclusive.

To investigate the reliability of consecutive-observation verification, an analytical false-clearance study was conducted using a stationary binary first-order Markov error model. The marginal probability of a joint false pass was denoted by (p), while temporal dependence between successive decisions was represented by the lag-one correlation coefficient (ρ). An absorbing dynamic-programming procedure was used to calculate the exact probability that at least one run of (N) consecutive false passes would occur within (H) eligible observations while unacceptable residual contamination persisted [10].

The principal analytical scenario used ($H=30$), with marginal false-pass probabilities of ($p=0.05$) and ($p=0.10$), correlation levels of ($\rho=0$) and ($\rho=0.70$), and confirmation requirements of ($N=1$), ($N=3$), and ($N=5$). The calculation was implemented in Python using double-precision arithmetic. Its numerical correctness was checked through independent enumeration of binary sequences for selected horizons and parameter combinations. These calculations represent analytical scenarios rather than measurements obtained from an actual camera or cleaning device.

Intervention record and observation eligibility

Each cleaning attempt receives an intervention identifier. Before dispatch, the monitor stores an immutable checkpoint containing the latest valid contamination estimates, their acquisition time, the camera configuration, and the applicable acceptance thresholds. The record also includes the command time, completion or fault indication, and the maximum duration allowed for verification. A completion indication establishes that the controller has finished its action; it supplies no direct proof of visual recovery [11].

After completion, the verifier waits for a settling interval selected for the actual hardware. This interval is intended to exclude spray, wiper motion, and short-lived optical or exposure disturbances. It is a configurable parameter rather than a universal constant. An observation is eligible only if its acquisition timestamp follows settling, its image identifier has not already been counted, its

configuration matches the episode, and all required measurements are valid. Processing time cannot substitute for acquisition time because a delayed result may describe a pre-cleaning image.

Observations and actuator messages must match the active intervention. A stale message from a previous cycle cannot open or close the current verification window. The verifier requests fresh measurements even when routine monitoring would otherwise reuse a previous mask. A change in crop, resolution, calibration, or critical-region definition ends the current comparison as inconclusive. A new assessment can then start with the new configuration.

Acceptance evidence

Let S denote the estimated fraction of the monitored image affected by contamination. Let C denote the maximum contaminated fraction among a fixed set of task-relevant image regions. Both quantities lie between zero and one and are supplied by the existing image-analysis subsystem. The regions and thresholds are frozen within an intervention. Their selection is application dependent; the present paper does not derive them or equate image coverage with accident risk.

Let U be a Boolean image-usability check with an independently specified validation target. In a controlled bench experiment, U could be based on a known optical target under fixed illumination. In operation, it requires a validated task-specific observable whose availability can be checked. Mean object-detector confidence is insufficient by itself because an empty scene can produce no detections without establishing image quality. When the required evidence cannot be obtained, U is unavailable and the observation is ineligible [12].

For an eligible observation, the joint pass decision is

$$B_t = \mathbf{1}\{S_t \leq \theta_S\} \mathbf{1}\{C_t \leq \theta_C\} U_t. \quad (1)$$

The thresholds express an application-specific acceptable residual condition, not a claim that the lens is perfectly clean. They must be selected using held-out intervention data. If U is unavailable, the system does not assign it a default value of one. Likewise, an unobserved region is not assigned zero contamination merely because no current mask was produced.

Bounded confirmation and terminal outcomes

The verifier accepts the episode when N eligible observations pass consecutively, with acquisition-time gaps no larger than a configured maximum. A valid failing observation resets the run length to zero. A scheduled observation with missing required evidence also breaks the run. Duplicate or unrelated messages are ignored and never extend it; an elapsed-time check resets a run if the maximum gap is exceeded. Counting follows acquisition order after a bounded reordering allowance. Results arriving after the verification deadline cannot retroactively change a terminal decision [13].

Two limits bound the procedure: a wall-clock deadline and a maximum number H of eligible observations. Whichever is reached first closes the episode. Acceptance is checked on the H -th observation before closing for the observation limit. A successful decision is recorded with the identifiers of the observations that support it. The terminal record is immutable; ordinary camera-health monitoring continues and can subsequently detect renewed contamination. Recovery verification is therefore time-local evidence rather than a permanent camera-health certificate.

Table 1. Episode outcomes and their interpretation

Outcome	Required interpretation
Verified acceptable	A complete qualifying run satisfies all acceptance checks within both limits.
Partial improvement	A comparable final window shows reduced contamination beyond tolerance, but residual acceptance conditions are not all met.
Unchanged	Comparable final measurements remain within the predeclared change tolerances.
Worsened	A comparable final window shows increased contamination beyond tolerance in either monitored quantity.
Inconclusive	Evidence is missing, incompatible, unstable, or insufficient to support one of the other outcomes.

For nonaccepted episodes, comparison uses the last N consecutive eligible observations satisfying the gap rule, if available. The medians of S and C form the final descriptor. Both descriptor ranges must stay within predeclared stability tolerances; otherwise the outcome is inconclusive. Relative to the checkpoint, a decrease larger than its change tolerance in at least one descriptor, with no increase beyond tolerance in either, gives partial improvement. An increase beyond tolerance in either descriptor takes precedence and gives worsened. Changes within both tolerances give unchanged. These rules leave reduced total area with increased critical-region obstruction visibly adverse rather than averaging the two effects.

A missing checkpoint prevents a claim of improvement or deterioration. It also makes the cleaning episode inconclusive under this procedure; a separate current-condition assessment may still be performed. An actuator fault is retained as an execution-fault flag and closes the verification attempt as inconclusive. Failure to verify recovery is not sufficient evidence of permanent lens damage. Additional inspection is needed to distinguish persistent soiling, internal condensation, a failed mechanism, and physical damage.

Retries require a new intervention record, an externally specified attempt limit, and a minimum interval between attempts. Resource and device constraints belong to the cleaner controller. The verifier reports its outcome and supporting measurements; it does not issue unrestricted repeated cleaning commands.

Analytical Study of False Clearance

Assumed error process

The numerical study isolates one remaining weakness after message and timing checks: the image-analysis subsystem can repeatedly pass an actually unacceptable image. Define a false-clear event as at least one run of N joint false passes while unacceptable residual contamination persists throughout the verification window. All H observations are assumed eligible and regularly spaced, so no gap or missing-evidence reset occurs. The analysis concerns this stated condition, not the complete system's field failure rate [14].

Let the joint decision B follow a stationary binary first-order Markov process. Its marginal false-pass probability is p and its lag-one correlation is ρ . For the nonnegative correlations considered here, transition probabilities are

$$a = P(B_t = 1 \mid B_{t-1} = 1) = p + \rho(1 - p), \quad b = P(B_t = 1 \mid B_{t-1} = 0) = p(1 - \rho). \quad (2)$$

The first observation is Bernoulli with probability p. Independence is recovered when ρ is zero. Under this model, the probability that one specified block of N observations all passes is p multiplied

by a to the power $N-1$. This block probability is not the probability of finding a qualifying run anywhere in H observations because candidate runs overlap.

Exact finite-horizon calculation

An absorbing dynamic program computes the probability over all candidate runs. After each observation, $v(r)$ stores the probability that acceptance has not occurred and the current run length is r , for r from zero to $N-1$. F stores accumulated acceptance probability. Initially, $v(0)$ is $1-p$; for N greater than one, $v(1)$ is p and F is zero. For N equal to one, F is p instead.

For each additional observation, start an empty next-state vector. From run length zero use pass probability b ; from a positive run length use a . Send failing probability mass to run length zero. Send passing mass to the next run length, except that a transition reaching N is added to F . Repeat until H observations have been processed. F is the requested false-clear probability. Probability mass is conserved at every step, and the computation requires $O(HN)$ operations and $O(N)$ storage.

The calculation was implemented in Python using ordinary double-precision arithmetic. Independent enumeration of all binary sequences checked horizons of 5, 8, and 10, run lengths of 1, 3, and 5, marginal probabilities of 0.05 and 0.10, and correlations of 0 and 0.70. Across 36 parameter combinations, the largest absolute discrepancy was below 2×10^{-15} . This checks the probability calculation for those cases; it does not validate the assumed error process against camera data.

Results and Interpretation

Table 2 reports the exact probabilities for $H = 30$. Every p and ρ value is an illustrative assumption. Values are percentages of hypothetical verification episodes under persistent unacceptable contamination, not percentages of incorrectly segmented pixels.

Table 2. Calculated false-clear probability within 30 eligible observations

p	ρ	$N = 1$	$N = 3$	$N = 5$
0.05	0.00	78.536%	0.333%	0.000773%
0.05	0.70	38.712%	20.658%	10.324%
0.10	0.00	95.761%	2.506%	0.023498%
0.10	0.70	62.793%	38.149%	20.754%

For $p = 0.10$ and $N = 3$, positive correlation increases the calculated probability from 2.506% to 38.149%. Increasing N to five lowers both values, but the correlated case remains 20.754%. A longer confirmation run is therefore not a substitute for measuring error persistence or improving the evidence used by the verifier.

The $N = 1$ rows show a different pattern: positive correlation reduces the probability of encountering any pass within the fixed window. Errors occur in longer clusters separated by longer failing periods. Once a cluster begins, however, it is much more likely to satisfy a multi-observation rule. This distinction explains why correlation cannot be summarized as a universal multiplier on false-clear probability.

Repeated opportunities also matter. With nonzero transition probabilities and no deadline, a finite run will eventually occur with probability approaching one, even if p is small. Limiting H prevents indefinite evidence accumulation, but it does not establish an acceptable application risk. That requires measured parameters, a justified risk target, and an assessment of model adequacy.

Confirmation has a latency cost even when all observations correctly pass. With the first eligible observation exactly at the end of settling and regular interval Δ , earliest confirmation occurs at settling duration plus $(N-1)\Delta$ after actuator completion. An illustrative 0.5 s settling time and 0.2 s

interval give 0.5 s, 0.9 s, and 1.3 s for N equal to one, three, and five. Processing delay, acquisition alignment, failed checks, and missing observations increase these times. These values are scheduling calculations, not measured execution times [15].

Worked Verification Cases

Table 3 applies the rules to constructed examples. They illustrate decision boundaries and are not a sampled dataset or an accuracy evaluation. Assume $N = 3$, an allowed gap of 0.25 s, and otherwise valid configuration and timing. A pass means all conditions in Equation (1) are met.

Table 3. Constructed cases illustrating verification behavior

Available evidence	Decision under the proposed procedure
Completion acknowledgment with no fresh images	Inconclusive at the deadline; completion alone cannot verify recovery.
Three passes acquired 0.2 s apart after settling	Verified acceptable if they fall within both episode limits.
One passing image delivered three times	One observation counted; no qualifying run.
Three passing results acquired before settling	Results excluded, even if processing completes later.
Two passes, missing required evidence, then one pass	Run restarts; only the last pass belongs to the new run.
Three passes with a 0.6 s gap before the third	Gap breaks the run; one current pass is retained.
Total affected area decreases while a critical region worsens	No acceptance if C exceeds its threshold; stable comparable evidence supports worsened.
Residual estimates pass but usability evidence is unavailable	Observation ineligible; no recovery claim.

The examples show why increasing N alone cannot address stale data, duplicated images, or an unavailable usability measure. Those conditions require explicit eligibility rules. Conversely, eligibility rules cannot remove a detector's systematic false negatives, which is the weakness quantified by the analytical model.

Experimental Validation Requirements and Limitations

A physical evaluation should record complete cleaning episodes, including a clean reference condition where feasible, controlled contamination, actuator telemetry, post-cleaning images, and an independent outcome assessment. The same episode must remain within one training, validation, or test partition. Randomly splitting its adjacent images would allow nearly identical contamination states to appear on both sides of the evaluation.

Thresholds, settling duration, permitted gaps, N, and H should be selected using validation episodes and fixed before testing. The main outcome is the proportion of truly unacceptable episodes incorrectly verified as acceptable. Complementary measures are the proportion of acceptable episodes verified before the deadline, time to verification, inconclusive-episode frequency, and cleaning attempts per episode. Reporting only frame-level segmentation accuracy would leave the intervention decision unevaluated.

The Markov study assumes stationary errors and memory of only the preceding decision. Actual errors may persist because of one unrecognized defect for an entire episode, change with exposure, or vary across contamination types. A fitted first-order model must therefore be checked against

observed run lengths and episode variability. Even a high-quality fit would apply only to the evaluated operating conditions. The calculation also excludes missing observations; its results must not be transferred unchanged to a different sampling policy.

Before-and-after improvement establishes a temporal association with the intervention. During vehicle motion, changing illumination or scene content can alter image-quality measurements independently of cleaning. Controlled targets or matched reference measurements are needed to isolate cleaning effects experimentally. In operational use, the system should record the verified post-action condition without claiming a causal estimate of cleaning effectiveness.

The practical dependency is the usability check U. Until an appropriate observable and its validation are supplied, the verifier remains a method specification with analytical support. Its conservative handling of missing evidence can also increase inconclusive outcomes and verification delay. These costs must be assessed together with false clearance rather than hidden by reporting successful episodes alone.

Results

The analytical evaluation demonstrated a substantial effect of temporal dependence on the reliability of consecutive-observation verification. With a verification horizon of 30 eligible observations, the probability of false clearance varied considerably according to the assumed marginal false-pass probability, temporal correlation, and required confirmation-run length.

For ($p=0.05$) under independence ($(\rho=0)$), the calculated false-clear probability was 78.536% when only one passing observation was required. Requiring three consecutive passes reduced this probability to 0.333%, while five consecutive passes reduced it further to approximately 0.000773%. However, when lag-one correlation increased to ($\rho=0.70$), the corresponding probabilities were 38.712%, 20.658%, and 10.324% for ($N=1$), ($N=3$), and ($N=5$), respectively.

A similar pattern was observed for the higher marginal false-pass probability of ($p=0.10$). Under independence, the probability of false clearance was 95.761% for ($N=1$), 2.506% for ($N=3$), and 0.023498% for ($N=5$). In contrast, under positive temporal correlation of ($\rho=0.70$), the probabilities were 62.793%, 38.149%, and 20.754%, respectively. Thus, for ($p=0.10$) and ($N=3$), introducing a lag-one correlation of 0.70 increased the calculated false-clear probability from 2.506% to 38.149%.

The results further showed that increasing the required number of consecutive passes reduced false-clear probability under both independent and correlated error conditions. Nevertheless, the reduction was considerably weaker when errors were temporally correlated. For example, increasing the requirement from three to five consecutive passes at ($p=0.10$) and ($\rho=0.70$) reduced false clearance from 38.149% to 20.754%, which remained substantially higher than the corresponding independent-error probability of 0.023498%.

The constructed verification cases additionally demonstrated the importance of observation eligibility. Three valid post-settling passes acquired within the permitted interval resulted in a verified-acceptable outcome, whereas repeated delivery of the same passing image counted as only one observation. Similarly, observations acquired before the settling period were excluded, missing usability evidence rendered an observation ineligible, and excessive acquisition-time gaps interrupted an otherwise successful confirmation sequence.

These findings indicate that successful completion of the mechanical cleaning action alone does not provide sufficient evidence of visual recovery. Verification requires fresh and temporally valid post-intervention observations satisfying both residual-contamination and image-usability criteria.

Discussion

The results demonstrate that persistence-based verification can strengthen confidence in camera recovery only when the temporal structure of decision errors is considered. Requiring several

consecutive acceptable observations appears highly effective under an independence assumption. However, this apparent reliability can deteriorate substantially when false-pass decisions occur in correlated clusters. The difference between 2.506% false clearance under independence and 38.149% at ($\rho=0.70$) for ($p=0.10$), ($N=3$), and ($H=30$) illustrates the practical importance of temporal dependence.

This finding is particularly relevant to intelligent vision systems because consecutive camera frames are rarely completely independent. Persistent contamination, stable lighting conditions, repeated segmentation errors, or an unrecognized optical defect can cause similar errors across multiple observations. Consequently, simply increasing the required number of consecutive passes cannot substitute for validating the underlying contamination detector and image-usability assessment. This is consistent with previous research showing that camera contamination can affect both image quality and object-detection performance [1], while confidence values produced by neural networks may not always be adequately calibrated [6].

The findings also support separating physical cleaning-system efficiency from image-based recovery verification. ISO 24650:2024 addresses the assessment of cleaning-system efficiency through contamination removal from the sensor front surface [2]. The procedure examined in the present study addresses a complementary question: whether sufficient post-intervention visual evidence exists to classify the camera condition as acceptable. Therefore, confirmation that an actuator successfully completed a cleaning command should not automatically be interpreted as confirmation that visual functionality has been restored.

Observation eligibility represents another important component of the proposed architecture. The worked cases demonstrate that stale images, duplicated observations, pre-settling frames, missing usability evidence, and excessive temporal gaps cannot be corrected merely by increasing (N). These cases require explicit intervention association, timestamp validation, configuration consistency, and evidence-availability rules. This distinction is particularly important for asynchronous intelligent systems, where the time at which an analysis result is received may differ substantially from the acquisition time of the underlying image.

The results further indicate that a finite observation horizon is necessary. Without a bounded horizon, repeated opportunities can eventually produce a qualifying sequence even when the marginal false-pass probability is relatively small. Limiting the number of eligible observations therefore prevents indefinite accumulation of evidence. Nevertheless, selecting appropriate values of (N), (H), settling duration, acceptance thresholds, and maximum observation gaps requires empirical validation rather than analytical assumptions alone.

Several limitations should consequently be acknowledged. The reported probabilities were generated from an assumed first-order Markov model and should not be interpreted as measured failure rates of an actual camera-cleaning system. No physical cleaner, real-world intervention dataset, or experimentally validated recovery classifier was evaluated in this study. Furthermore, actual decision errors may exhibit higher-order temporal dependence, contamination-specific behavior, exposure-dependent variation, or episode-level persistence that cannot be completely represented by a stationary first-order Markov process.

Future experimental research should therefore evaluate the proposed verification procedure using instrumented cleaning episodes containing controlled contamination, actuator telemetry, pre- and post-cleaning images, and independently established recovery labels. Such experiments should measure episode-level false-clear rates, successful verification rates, verification latency, inconclusive outcomes, and the number of cleaning attempts required. These measurements would provide the empirical evidence necessary to determine operational thresholds and assess whether the proposed analytical verification architecture translates into reliable real-world camera-health monitoring.

Conclusion

The proposed procedure verifies camera-cleaning outcomes using observations tied to one intervention, explicit settling and freshness requirements, a joint residual-contamination and usability check, and bounded consecutive confirmation. It distinguishes acceptable post-action evidence from completion of a cleaning command and preserves inconclusive outcomes when the evidence is insufficient.

Exact calculations under an assumed Markov error model show that temporal dependence can greatly weaken the apparent benefit of consecutive passes. For $p = 0.10$, three passes within 30 observations produce 2.506% false clearance under independence and 38.149% at correlation 0.70. The result supports measuring decision-error persistence and setting a finite observation horizon. Validation with instrumented cleaning episodes is necessary before selecting operational thresholds or making claims about real-camera reliability.

References

- [1] H. Kim, Y. Yang, Y. Kim, D.-W. Jang, D. Choi, K. Park, S. Chung, and D. Kim, “Effect of droplet contamination on camera lens surfaces: Degradation of image quality and object detection performance,” *Applied Sciences*, vol. 15, no. 5, Art. no. 2690, 2025, doi: 10.3390/app15052690.
- [2] International Organization for Standardization, *ISO 24650:2024—Road Vehicles—Sensors for Automated Driving Under Adverse Weather Conditions—Assessment of the Cleaning System Efficiency*. Geneva, Switzerland: ISO, 2024.
- [3] A. G. O. Axmedov, “AI-based image quality and lens defect analysis in autonomous driving: A framework with U-Net-based soiling detection,” *American Journal of Mechanics and Applications*, vol. 12, no. 4, pp. 93–101, 2025, doi: 10.11648/j.ajma.20251204.14.
- [4] A. Das, P. Krizek, G. Sistu, F. Burger, S. Madasamy, M. Uricar, V. R. Kumar, and S. Yogamani, “TiledSoilingNet: Tile-level soiling detection on automotive surround-view cameras using coverage metric,” in *Proc. 2020 IEEE 23rd Int. Conf. Intelligent Transportation Systems (ITSC)*, 2020, pp. 1–6, doi: 10.1109/ITSC45102.2020.9294677.
- [5] M. Uricar, J. Ulicny, G. Sistu, H. Rashed, P. Krizek, D. Hurych, A. Vobecky, and S. Yogamani, “Desoiling dataset: Restoring soiled areas on automotive fisheye cameras,” in *Proc. IEEE/CVF Int. Conf. Computer Vision Workshops (ICCVW)*, 2019.
- [6] C. Guo, G. Pleiss, Y. Sun, and K. Q. Weinberger, “On calibration of modern neural networks,” in *Proc. 34th Int. Conf. Machine Learning (ICML)*, vol. 70, 2017, pp. 1321–1330.
- [7] International Organization for Standardization, *ISO 21448:2022—Road Vehicles—Safety of the Intended Functionality*. Geneva, Switzerland: ISO, 2022.
- [8] M. Bijelic, T. Gruber, F. Mannan, F. Kraus, W. Ritter, K. Dietmayer, and F. Heide, “Seeing through fog without seeing fog: Deep multimodal sensor fusion in unseen adverse weather,” in *Proc. IEEE/CVF Conf. Computer Vision and Pattern Recognition (CVPR)*, 2020, pp. 11682–11692.
- [9] M. Hahner, C. Sakaridis, D. Dai, and L. Van Gool, “Fog simulation on real LiDAR point clouds for 3D object detection in adverse weather,” in *Proc. IEEE/CVF Int. Conf. Computer Vision (ICCV)*, 2021, pp. 15283–15292.
- [10] M. Cordts *et al.*, “The Cityscapes dataset for semantic urban scene understanding,” in *Proc. IEEE Conf. Computer Vision and Pattern Recognition (CVPR)*, 2016, pp. 3213–3223, doi: 10.1109/CVPR.2016.350.
- [11] H. Caesar *et al.*, “nuScenes: A multimodal dataset for autonomous driving,” in *Proc. IEEE/CVF Conf. Computer Vision and Pattern Recognition (CVPR)*, 2020, pp. 11621–11631.
- [12] A. Geiger, P. Lenz, and R. Urtasun, “Are we ready for autonomous driving? The KITTI vision benchmark suite,” in *Proc. IEEE Conf. Computer Vision and Pattern Recognition (CVPR)*, 2012, pp. 3354–3361, doi: 10.1109/CVPR.2012.6248074.
- [13] Z. Wang, A. C. Bovik, H. R. Sheikh, and E. P. Simoncelli, “Image quality assessment: From error visibility to structural similarity,” *IEEE Transactions on Image Processing*, vol. 13, no. 4,

- pp. 600–612, Apr. 2004, doi: 10.1109/TIP.2003.819861.
- [14] A. Kendall and Y. Gal, “What uncertainties do we need in Bayesian deep learning for computer vision?” in *Advances in Neural Information Processing Systems*, vol. 30, 2017.
 - [15] D. Hendrycks and T. Dietterich, “Benchmarking neural network robustness to common corruptions and perturbations,” in *Proc. Int. Conf. Learning Representations (ICLR)*, 2019.