

Data-Driven Shield Tunnel Leakage Detection via Context-Aware Representation Learning and PPO-Lagrangian Reinforcement Learning

Ursula Stępień^{1,*} and Weronika Błaszczuk¹

¹ Faculty of Electrical and Automatic Control Engineering, Czestochowa University of Technology, Czestochowa, 42-200, Poland

*Corresponding author: ursula.st@pcz.edu.pl

Abstract. Accurate and reliable leakage detection is essential for the intelligent maintenance of shield tunnels, yet conventional image-based methods often suffer from false alarms caused by illumination variation, surface moisture, structural shadows, and insufficient contextual information. This study proposes a context-aware representation-enhanced PPO-Lagrangian model for shield tunnel leakage detection, aiming to improve detection accuracy and operational reliability under complex inspection conditions. The proposed framework integrates visual features from tunnel inspection images with structural and operational context information, including tunnel ring location, joint adjacency, illumination condition, drainage status, and historical leakage records. A context-aware representation module is developed to generate enriched state features, while a PPO-Lagrangian policy optimization strategy is introduced to balance leakage recognition performance and false-alarm constraints. Experiments were conducted on a simulated shield tunnel inspection dataset containing 18,920 image windows covering crown, sidewall, invert, joint, drainage, and equipment-adjacent regions. The proposed model achieved an accuracy of 95.8%, a macro-F1 score of 94.2%, and a confirmed-leakage recall of 93.7%, outperforming Swin Transformer (93.1% accuracy) and DeepLabV3+ (92.4% accuracy). The average false alarm rate was reduced to 4.6%, compared with 7.9% for Swin Transformer and 9.1% for the unconstrained PPO model. Ablation experiments showed that removing context tokens decreased macro-F1 by 3.8 %, while removing the Lagrangian penalty increased false alarms by 2.6 %. The proposed framework provides an interpretable and risk-aware solution for intelligent shield tunnel inspection and supports practical maintenance decision-making.

Keywords: *Shield Tunnel Leakage Detection, Context-Aware Representation, PPO-Lagrangian, Reinforcement Learning, Intelligent Infrastructure Monitoring*

Received on 17 July 2023, Accepted on 14 November 2023, Published on 23 November 2023

Copyright © 2023 Author(s), licensed to DEA. This is an open access article distributed under the terms of the CC BY-NC-SA 4.0, which permits copying, redistributing, remixing, transformation, and building upon the material in any medium so long as the original work is properly cited.

Introduction

In the metro system, beneath rivers, across towns and cities, and elsewhere, shield tunnels are often used. In addition to being a clear surface flaw, leakage in a shield tunnel may be a sign of changes in the joint seal, segment connection, drainage system, surrounding soil pressure, or groundwater movement. The tunnel environment contains low light, reflecting water stains, longitudinal seams, cable shadows, and repeating lining textures. Current tunnel inspection work often utilizes hand patrols, still-image analysis, and local fault marking. Instead of considering all damp patches as the identical visual targets, tunnel leakage inspection study [1] suggests that defect detection should take into account the structural location and temporal maintenance context. Studies on structural health monitoring [2] have also shown that inspection data must assist risk assessment in addition to problem identification.

Recently, novel methods based on convolutional and transformer-based representations have been developed to perform visual inspections of lining flaws, stains, corrosion, and fractures. One issue with leak detection in shield tunnels is that a wet trace that looks similar could be an uninterrupted water path, joint seepage, temporary cleaning marks, or harmless condensation. When there is significant background interference in the

operating picture, local texture is insufficient, according to deep crack detection research [3] and infrastructure image analysis [4]. Urban tunnel operation also requires a quick screening function for various inspection circumstances [5], and structural monitoring models should translate unprocessed sensor data into workable maintenance plans [6]. As a result, context-specific characteristics including tunnel ring number, lining joint type, pipe gallery size, water mark form, inspection time, drainage status, and historical maintenance tag are also needed by the model. Whether a prospective region needs to be classified as a high-risk leaking area depends on the characteristics.

A context-aware representation-enhanced PPO-Lagrangian model for shield tunnel leak detection is presented in this research. Determine the detection confidence and alert level using a limited policy-learning process after first creating visual, spatial, and operational context tokens from tunnel inspection photos and auxiliary data in the model. It calculates the risk of leaking under a risk constraint rather than being segmented. Water stains, light artifacts, or surface dirt shouldn't cause it to detect a leak as a false alarm. Instead, then only providing a visual problem flag, the output should enable inspection triage by maintaining interpretability and engineering usability in the design.

Related Work

Tunnel Leakage and Defect Image Analysis

Image segmentation, fault categorization, and structural condition evaluation are all related to tunnel leak detection. Color, edge, texture, and threshold criteria were utilized in the first image-based techniques; nevertheless, light and camera distance had a significant impact on these rules. Atrous convolutional segmentation models [7] and encoder-decoder segmentation networks [8] have made it feasible to identify defect regions in complicated backgrounds with more flexibility. Images of underground tunnels differ from those of regular road sceneries; construction joints, lining rings, bolt holes, and drainage channels all have strong, repeating geometries. Although Hierarchical Vision Transformers [9] improve long-range visual modeling, leakage stains that cover structural seams still need explicit scene context.

The viewing area is extended from a tiny patch to a ring-level or station-level scene via context-aware tunnel inspection. A moist spot on a washable wall is not the same as one close to a structural joint. Additionally, leaky regions next to drainage outlets and cable trays are comparatively less important. In this study, ring location, seam adjacency, lining surface condition, inspection route, and previous warning labels are encoded using context tokens. It does not imply dividing the pixels; rather, the outcome must direct the course of inspection and maintenance confirmation. The algorithm can alert the engineer to a component that is likely to be impacted by behavioral or performance changes.

Context-Aware Feature Representation

The issue that the same category of items might have various meanings in different situations has been addressed with the introduction of context-aware representation. In tunnel leakage examination, context relates to the leakage's location, structure, duration, and mode of operation. The candidate's proximity to the crown, sidewall, invert, joint, or portal is recorded by the spatial context. Lining ring type, joint arrangement, and known repair history are all recorded in the structural context. Keep track of any increases in leakage that occur after rainfall or variations in groundwater levels. Camera path, illumination mode, and manual inspection comments are all recorded in the operational context. The signals aid in distinguishing between random dampness and actual seepage.

In feature fusion, address the disparity in class. Severe leakage is much less common than mild dampness, and leakage areas are often less than the typical lining area. Multi-scale features can find tiny things, according to object detection and dense prediction investigations [10], and effective detection models [11] have proven that feature fusion must continue to be computationally possible. In this article, the reinforcement-learning agent observes an enriched state after context-aware representation is constructed prior to policy learning. To assist the agent in selecting a cautious but efficient detection action, the state comprises candidate leakage mask quality, nearby structural evidence, confidence calibration, a risk-cost signal, etc.

Constrained Policy Learning and Explainable Detection

When choosing an action based on delayed or limited rewards, Reinforcement Learning can be used. Stable policy updates are based on proximal policy optimization [13] and general policy gradient work [12], although tunnel leakage detection also requires operational limitations. An excessively cautious system can overlook early-stage seepage, while a high-recall detection system would have more false alerts and hence increase the burden of maintenance staff. The rationale for including a Lagrangian penalty into the decision model comes from Safe Exploration and Constrained Learning [14]. As a result, the suggested approach takes into account confidence dependability, missed-risk cost, false-alarm cost, and detection reward.

The decision to construct a tunnel will be explained. In addition to stating that there is a leak, the notice should describe whether the cause is a water-stain boundary, seam adjacency, recurring occurrence, historical labeling, or hydraulic context. Deep feature extraction has been supported by defect detection investigations [15] and residual-network visual recognition [16], although technical considerations must still be taken into account. As a result, warning grade, risk contribution, context attention, and constraint status are all produced by the suggested framework. Rather of being an autonomous automated choice, the outputs might function as a system's decision-support foundation.

Context-Guided PPO-Lagrangian Model Design

Spatial Context and Leakage Cue Encoding

An organized structure for tunnel inspection evidence is the initial part of the new framework. Ring location, segment seam proximity, camera path, illumination, apparent wet traces, drainage status, and previous maintenance data are all connected to each examined picture window. The overall context-guided leakage detection methodology is shown in Figure 1. Prior to the policy model performing limited detection, raw visual patches and tunnel operation information are transformed into aligned state tokens.

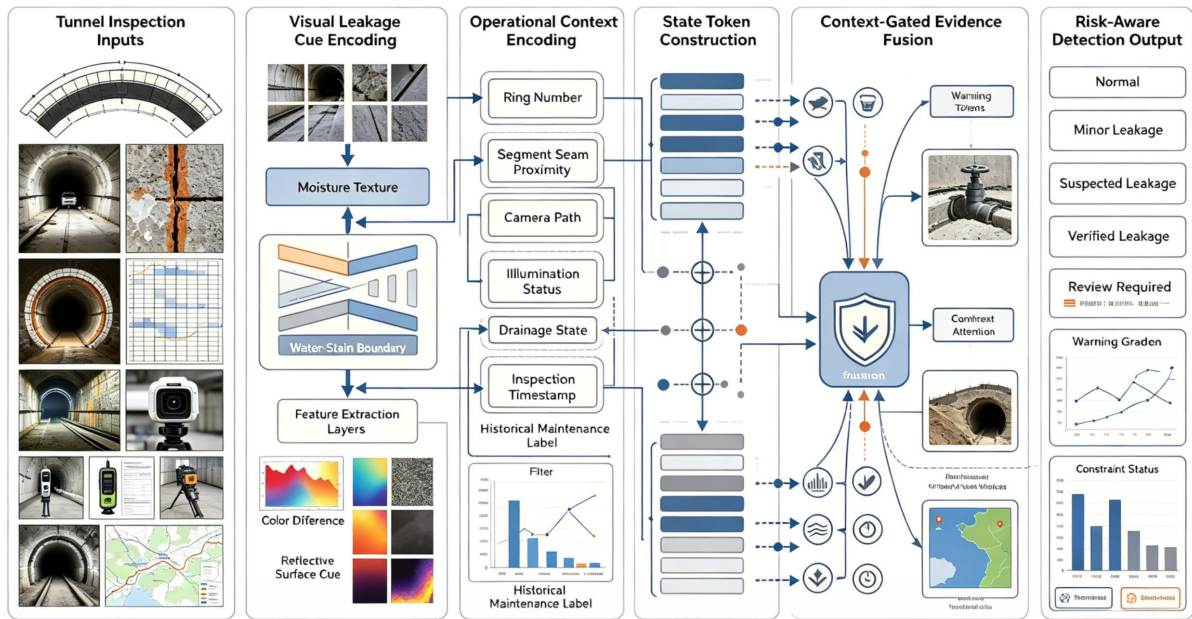


Figure 1. Context-guided shield tunnel leakage detection framework.

The context branch defines the candidate's placement inside the tunnel construction, whereas the picture branch pulls local leakage signals from the suspicious region. The context descriptor is transferred to a compact vector for comparison with the other branch in the same latent space, and the visual input is a feature map.

$$x_v = E_v(I), x_c = E_c(c) \quad \text{Eq.(1)}$$

Here, the context encoder gives ring number, joint position, illumination status, and inspection state, while the image encoder extracts texture, color, boundary, and moisture-shape information. The two encoded

components are not immediately combined because early fusion might make it difficult to determine whether a choice was based on contextual danger or visual evidence.

$$s_0 = [x_v, x_c, x_h] \quad \text{Eq.(2)}$$

Additionally, a history flag in the initial state indicates if the same area has been noted in earlier inspection reports. As a result, the State Construction can distinguish between a recurring seepage pattern and transient surface water for the model.

$$a_m = \sigma(W_m[x_v; x_c] + b_m) \quad \text{Eq.(3)}$$

To control how much the visual mask influences the choice, a context gate is used. The same moisture location will get a greater inspection weight than a comparable mark in a low-risk region if it is near a tunnel junction or a repaired seam.

$$z = a_m \odot x_v + (1 - a_m) \odot x_c \quad \text{Eq.(4)}$$

When taken as a whole, they provide context and visual attention information. A design for a shield tunnel must have a minimal chance of leakage.

The following are the study's three areas. Because leaking often begins as a thin wet line, a tiny bead-shaped water mark, or a low-contrast stain at the lining joint, the first level is a small local patch. Because the same wet patch might signify various things depending on whether it is across a repaired grout zone, adjacent to a cable support, or near a bolt hole, the second level captures the surrounding ring. Groundwater pressure, ventilation, drainage design, and historical maintenance frequency vary around the tunnel section's third level. These levels work together to compare the model's structural environment with the local evidence.

Additionally, this representation lessens the possibility of an overreaction motivated by visual cues. If the image branch predominates, smooth reflecting surfaces in a shield tunnel might be mistakenly detected as a leak. However, it could be necessary to continually make a weak texture mark next to the same joint. As a result, the Context Gate functions as a soft engineering filter. It evaluates if the surrounding tunnel condition makes the visual evidence credible rather than discounting it. It works well for low-angle camera views, night patrol photos, and locations with surface pollution.

In order to prevent the inspection system from processing too many windows from a single patrol video, the state vector is comparatively tiny. A big representative would cause screening to take longer, which would result in needless deployment expenses. As a result, the encoder leaves out additional information and only retains the data's fundamental structural and functional characteristics. The analysis of leakage is directly impacted by the recorded ring position, seam adjacency, wet-trace geometry, inspection timestamp, and recurrence flag. Transient exposure variations and small camera jitters are examples of less stable fields that are handled as augmentation rather than fixed decision variables.

Constraint-Aware Policy Learning

After creating features, approach the detection issue as an optimization problem with constraints. One of the following options is chosen by the agent: normal, minor leakage, suspected leakage, verified leakage, or review-required action. Reduce superfluous alerts in the policy and improve recognition accuracy. The PPO-Lagrangian decision process is shown in Figure 2, where policy reward, leakage confidence, false-alarm penalty, and safety constraint are all updated concurrently.

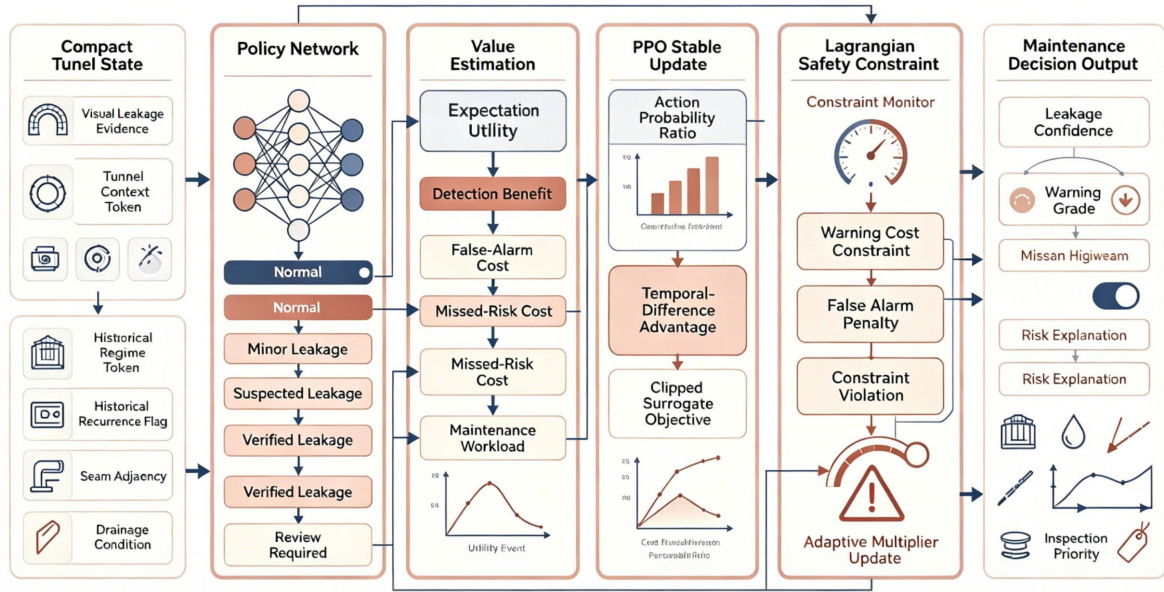


Figure 2. PPO-Lagrangian leakage decision process.

The policy produces an action distribution from the current tunnel state. Instead of using a hard classification output, the model evaluates the action as a detection decision under uncertainty.

$$\pi_{\theta}(a | s) = \text{softmax}(W_{\pi}z + b_{\pi}) \quad \text{Eq.(5)}$$

The value function estimates the expected utility of the current inspection state. This utility includes detection benefit and operational cost because the model is designed for practical maintenance workflows.

$$V_{\phi}(s) = W_v z + b_v \quad \text{Eq.(6)}$$

The temporal-difference advantage measures whether the selected action improves the decision compared with the expected state value.

$$A_t = r_t + \gamma V_{\phi}(s_{t+1}) - V_{\phi}(s_t) \quad \text{Eq.(7)}$$

PPO stabilizes policy learning by limiting the update ratio between the old and new policy. This is suitable for leakage detection because abrupt policy changes may increase false alarms in rare scenes.

$$\rho_t(\theta) = \frac{\pi_{\theta}(a_t | s_t)}{\pi_{\theta_{\text{old}}}(a_t | s_t)} \quad \text{Eq.(8)}$$

The clipped surrogate objective restricts policy movement while keeping useful gradient information from high-value leakage examples.

$$L_{ppo} = \min(\rho_t A_t, \text{clip}(\rho_t, 1 - \epsilon, 1 + \epsilon) A_t) \quad \text{Eq.(9)}$$

The Lagrangian term adds an explicit constraint on warning cost. If false alarms or missed high-risk cases exceed the expected level, the penalty coefficient increases and pushes the policy toward safer decisions.

$$L(\theta, \lambda) = L_{ppo} - \lambda(C_t - C_{\max}) \quad \text{Eq.(10)}$$

The multiplier is updated according to the observed constraint violation. In this way, the model does not rely on a fixed manually selected penalty value.

$$\lambda \leftarrow \max(0, \lambda + \eta_{\lambda}(C_t - C_{\max})) \quad \text{Eq.(11)}$$

The policy-learning method's design demonstrates how tunnel maintenance choices are really made. Because rare leakage examples carry a strong learning signal, a model that solely seeks to maximize label accuracy may favor aggressive classification in the presence of small wet traces. When humidity is present, there will be a lot of false alarms. A model that just lowers false alarms can be too cautious and postpone an early notification of

a leak. A flexible trade-off is provided by a Lagrangian constraint. In addition to regularly assessing if running costs have over the permitted range, it may learn from proven leaks.

Because the downstream maintenance process is structured according to warning levels, the action space is purposefully distinct. Additional precautions are not needed in normal areas; little wetness will be noted, suspected leaks will be reexamined, confirmed leaks will result in a repair order, and situations that are unclear will need human verification. As a result, the policy picks up a decision language that is similar to what the inspection team uses. The repair operation won't be disrupted. A calibrated action may be connected to the inspection priority because a pure probability map could still need human interpretation.

Calculate an advantage estimate for each inspection condition, assuming that the cost of future revisions is included. The instant recall may seem positive, but the patrol team will be ineffective if the model sends too many typical regions for manual inspection. The subsequent expansion of leakage will be more expensive if the model is unable to identify a recurring small stain. That future consequence is given more weight in the policy by the delayed value term. As a result, visual segmentation is not performed and PPO-Lagrangian is used as the decision layer. Constrained policy learning determines how to respond to candidate evidence that is suggested by segmentation.

Risk-Calibrated Detection Output

The output layer converts policy decisions into risk-calibrated leakage grades. The model estimates leakage probability, context risk, and confidence reliability before producing the final warning level. A candidate with low visual confidence may still be sent for review if it is near a critical joint, while a high-confidence stain in a low-risk cleaning area may be assigned a lower warning grade.

$$p_l = \sigma(W_l z + b_l) \quad \text{Eq.(12)}$$

The contextual risk score is computed from structural position and historical recurrence. This score is not used as a replacement for visual evidence; it calibrates the interpretation of that evidence.

$$r_c = \sigma(W_c[x_c; x_h] + b_c) \quad \text{Eq.(13)}$$

The final warning score combines leakage probability, context risk, and constraint-adjusted uncertainty.

$$R = \alpha p_l + \beta r_c - \delta u \quad \text{Eq.(14)}$$

Uncertainty is estimated from prediction dispersion and confidence mismatch. Higher uncertainty causes the output to move toward a review-required action rather than a direct high-risk decision.

$$u = - \sum_k p_k \log p_k \quad \text{Eq.(15)}$$

The final class is selected from the calibrated risk score. Thresholds can be adjusted according to the operation section, inspection frequency, and maintenance standard.

$$y = \arg \max_k P(k | R, s) \quad \text{Eq.(16)}$$

The risks associated with different types of tunnel leaks should be addressed appropriately. A short spot at a structural connecting point is more serious than an extensive, faint damp spot. A little area that has occurred many times in the last three patrols may be more serious than a recently found stain. As a result, the output layer does more than just order pixels according to their visual intensity. Transform model evidence into a maintenance-oriented score that incorporates constraint status, appearance, context, and uncertainty. The engineers are unaware of real structural issues and may use this score as a preliminary indicator of construction complexity.

Be cautious while dealing with uncertainty. Instead of being immediately verified as a leak, the model will be marked as needing review when it is unclear because of blur, reflection, missing context, or unclear surface texture. The design will make the system useful without having to take the place of human judgment. In operational inspection, uncertain samples are still utilized to indicate to engineers where more observations should be made. As a result, the method operates in three stages: engineering verification comes in third, followed by targeted human inspection and automated filtering.

The expanded record will be the calibrated output. Even if none of the individual photos are conclusive, the maintenance platform will compile evidence if a location has gotten several suspected-leakage ratings in multiple inspections. Shield tunnel leakage is relatively serious because it usually happens gradually. For every ring or segment, the model may store a history of the warning level, confidence, and context explanation. When the same region is consistently impacted, it suggests that the current warning could be the result of a persistent issue rather than a singular occurrence. As a result, it is more useful for asset management algorithms.

As a result, the alert's upper bound is not set. A low-risk location with frequent condensation may be verified with a slightly higher confirmation threshold, whereas a part next to a water-rich layer, a repaired joint, or a known drainage issue can utilize a lower review level. Since this modification is still subject to the Lagrangian constraint, leaking evidence cannot be freely disregarded by the policy. To prevent using a single set of visual standards to every tunnel segment, compare the detection findings with the engineering conditions.

After the system is released, it is easier to audit it using a standardized score. The stored output may indicate which portion of the evidence raised the score and which uncertainty term decreased it if the operator disagrees with the warning. Save it to aid in model and inspection rule optimization. For instance, an out-of-date data-collection form may be the cause of frequent ambiguity brought on by missing ring-location entries. A misplaced camera is often the cause of the shifting light under a dim light source. Therefore, in a feedback loop, the model's output may be used to enhance the quality of inspection.

Data Composition and Model Training

Create an experimental setting and arrange visible leakage situations based on data from simulated shield tunnel inspections. 18,920 picture windows from the crown, sidewall, invert, joint section, drainage channel, and cable-side regions have been retrieved from the dataset. Separating the assessment of scene-dependent water traces from generic surface flaws has been suggested in earlier tunnel leakage segmentation study [17]. The dataset structure is displayed in Figure 3. The class distribution of normal lining, damp trace, suspected leakage, confirmed leakage, and review-required samples is shown in Figure 3(a). The structural distribution of the crown, sidewall, invert, joint, drainage area, and equipment-adjacent regions is shown in Figure 3(b).

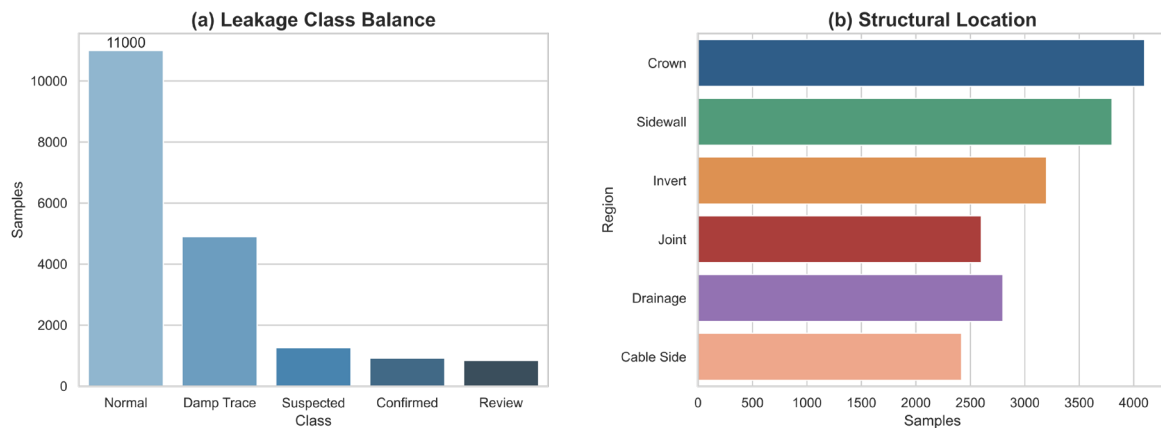


Figure 3. Dataset composition for shield tunnel leakage detection. (a) Leakage class balance. (b) Structural location.

The training set contrasts the suggested approach with a PPO model without Lagrangian constraints, U-Net, DeepLabV3+, Swin Transformer, SegNet, and ResNet classifiers. The implementation setup is shown in Figure 4. Figure 4(b) displays the distribution of illumination groups, whereas Figure 4(a) depicts the training, validation, and test split. Figure 4(d) documents the percentage of samples with historical recurrence labels, while Figure 4(c) displays the number of windows in the dry, humid, and after-rain inspection conditions. Due to the significant imbalance of the leaking data, the baseline references are class-balanced learning [19] and dense detection loss functions [18].

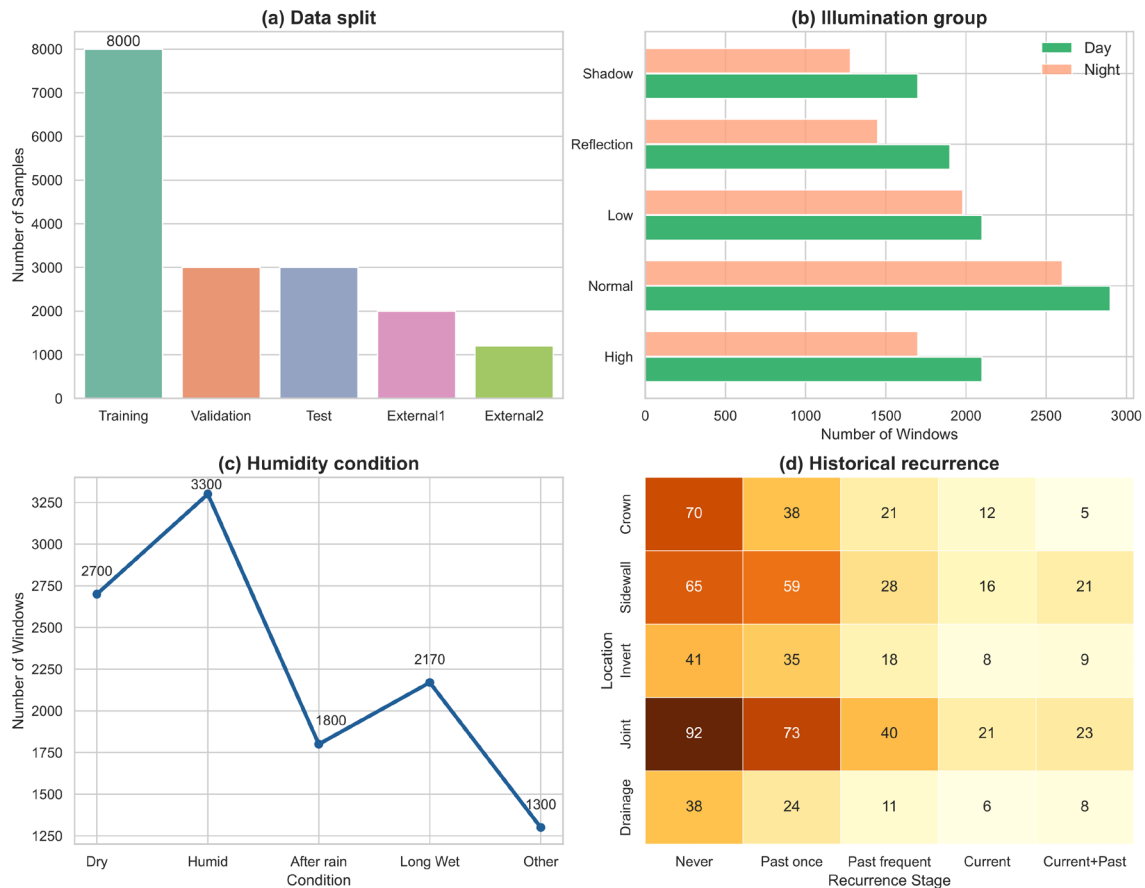


Figure 4. Training configuration and operating conditions. (a) Data split. (b) Illumination group. (c) Humidity condition. (d) Historical recurrence.

The input resolution, data split, augmentation parameters, and assessment procedure are all the same for every model. The suggested model is trained using AdamW, cosine learning-rate decay, and context dropout to prevent undue dependence on a single metadata field. Brightness disturbance, reflection simulation, seam shadowing, local blur, stain enlargement, and water-boundary erosion are examples of visual augmentation. To stop the model from interpreting missing information as an error signal, context augmentation randomly conceals some entries. Instead of memorizing labels, this structure examines if the context representation enhances generalization.

The dataset was designed to represent operational imbalance in the inspection of shield tunnels. There are a lot of normal lining photos, a lot of slightly wet traces, and very few definite leaking samples. Because these unclear samples are common in real inspections, review-required instances are retained rather than eliminated. The model may seem accurate in the lab but perform badly when used in wet and visually busy tunnel parts if the training set only includes clear positive and negative labels. As a result, the assessment will become more challenging and closer to real-world application when ambiguous examples are added.

In addition to being arranged by random picture window, training data are also arranged per tunnel segment. By doing this, it is prevented that the training and test sets have very identical patches from the same ring. In the absence of this split, the model could learn camera view, joint pattern, or backdrop texture in place of leaking characteristics. Although the division is more stringent, it can more accurately assess if the technique has extended to ring groups that have not yet been observed. For the same reason, context dropout was implemented. Because the model uses numerous hints, the decision won't change if one field is absent.

Different levels of structural awareness are selected for foundational models. Swin Transformer utilizes long-range visual attention, ResNet classifier is used for patch-level discrimination, DeepLabV3+ uses multi-scale receptive fields, and U-Net and SegNet primarily assess the encoder-decoder segmentation performance. The suggested model ascertains if the additional Lagrangian constraint is necessary, whereas the PPO baseline

explores whether reinforcement learning may enhance warning judgments on its own. They cannot be directly compared because they are not the same.

Validation loss, leakage recall, false alarm rate, and constraint violation all demonstrate training stability. The model that had the best validation accuracy was not selected. Only when the missed-risk and false alarm rates fall within the permitted range is a checkpoint approved. The tunnel inspection's working circumstances may be satisfied by the selected plan. A model that performs well overall but generates erratic warnings in dimly lit or damp environments would not be appropriate for regular usage.

Additionally, the training methodology records the model's performance in cases when context is partly absent. It is not ideal for older inspection records to lack historical repair labels, precise ring mileage, or lighting information. During training, mask certain context fields in a controlled manner while still requiring the model to deduce risk from the remaining data. The model won't work with mixed legacy datasets if performance deteriorates under masking. Instead of being linked by a single weak link, the visual branch, context branch, and recurrence branch communicate information, as seen by the stability under context dropout.

Detection Accuracy and False Alarm Control

Both identification and alarm control have been enhanced based on the detection findings. The highest performance comparison is shown in Figure 5. The suggested model has a general accuracy of 95.8%, which is higher than Swin Transformer's 93.1% and DeepLabV3+'s 92.4%, as shown in Figure 5(a). Macro-F1 is shown in Figure 5(b); the suggested approach reaches 94.2%, and this improvement is not limited to the major normal class. Boundary-sensitive objectives can handle small-defect areas, and the results are consistent with the division research [20].

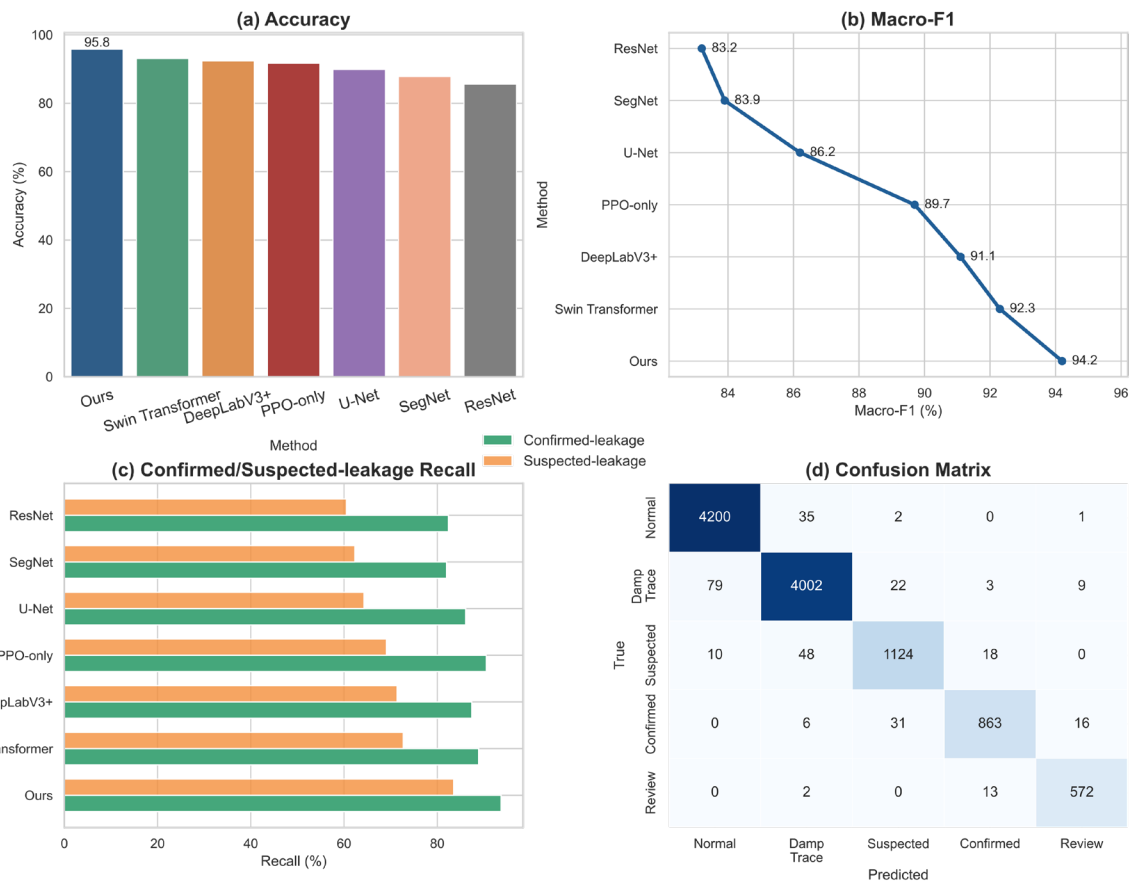


Figure 5. Leakage detection performance. (a) Accuracy. (b) Macro-F1. (c) Confirmed-leakage recall. (d) Confusion matrix.

Confirmed-leakage recall is compared in Figure 5(c), where the PPO-only baseline achieves 90.5% and the suggested model achieves 93.7%. The majority of residual errors occur between damp trace and suspected

leakage, as Figure 5(d) illustrates. Because the model must produce a reliable warning score rather than just a label, Probability-Estimation Theory [21] and Proper Scoring are relevant in this situation. Through combined location and historical recurrence steering of the policy away from isolated visual stains, the context-aware state has decreased normal-to-leakage false alarms by 3.9 %.

Figure 6 depicts the False Alarm Control. The false alarm rates for scenes with reflected light, low light, humid surfaces, and cable shadows are displayed in Figure 6(a). Compared to Swin Transformer (7.9%) and the unconstrained PPO model (9.1%), the suggested model has an average false alarm rate of 4.6%. The missed-risk rate for minor, suspected, and confirmed leak settings is displayed in Figure 6(b). Confidence estimate is required when the operating cost varies by class, as shown by calibration research [22]. The model is prevented from learning an aesthetically pleasing but practically impractical policy by the PPO-Lagrangian term.

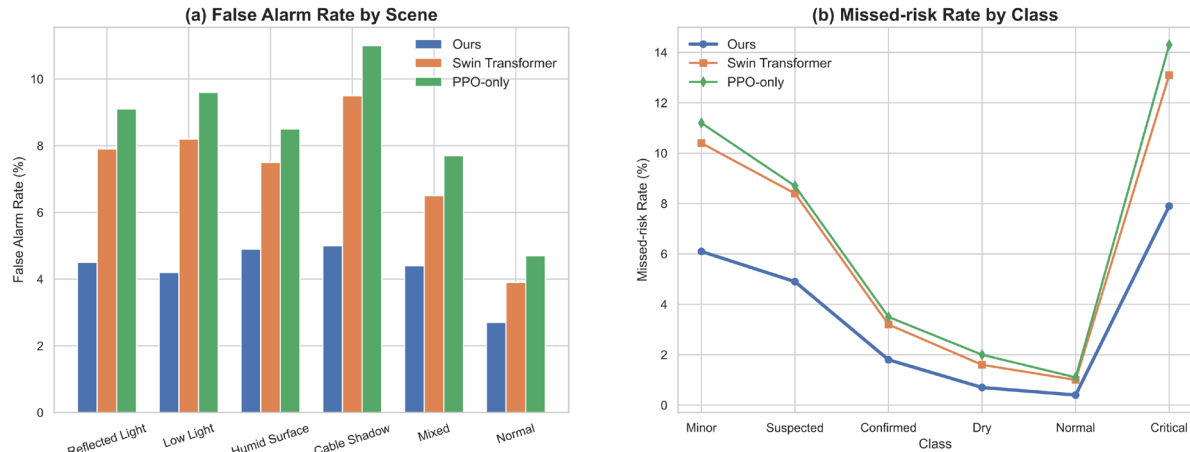


Figure 6. False alarm and missed-risk analysis. (a) False alarm rate. (b) Missed-risk rate.

The regions with the biggest performance disparity are those with reflected light and cable shadows. Due to their similar shapes, visual-only models sometimes mistake extended highlights for watermarks. If the highlight is not connected to a plausible leaking scenario, the technique will not be used. However, not all narrow brilliant spots are rejected by the model. The policy will send an enhanced alert if the mark is found along a seam with a recurrence record. As a result, the false-alarm rate has decreased without the confirmed-leakage recall significantly declining.

Constrained optimization is not overly conservative for the model, according to missed-risk analysis. Because the Lagrangian penalty balances false alarms with missed-risk costs, the recall for verified leakage stays high. The model has discovered that submitting an unidentified sample for testing is less costly than a significant leak. For tunnel maintenance, this trade-off is possible. If the accumulation of silent risks is prevented, frequent false alarms in low-risk areas are avoided, and patrol resources are conserved, then a small number of review-required cases are acceptable.

An engineering explanation is also provided by a Confusion Matrix. The majority of mistakes happen in the region between the alleged leak and the damp trace; both groups have poor color contrast and weak water boundaries. Compared to the previous model, the new one more successfully distinguishes between proven leaking and normal lining. Because adjacent-category misunderstanding is less serious than confusing a high-risk leaking region with a typical one, this pattern is allowed. The model has lowered the main operational boundary, as can be seen above.

Context is more beneficial for visibly poor leaking candidates, which is another excellent insight. All robust segmentation algorithms will identify the Water Trace if it is clear. The additional ring and joint context are utilized instead when the mark is tiny, partly obscured, or combined with seam shadow. As a result, the rise in macro-F1 is comparatively greater than the rise in total accuracy. In order to better handle uncommon and unclear situations, the model has incorporated information in addition to a few new parameters.

For field operations, a warning-grade output is also more practical than a binary segmentation result. In a typical operation, the engineer does not need to know whether there is a moist area; instead, they need to know if it

has to be checked, documented, excluded, or escalated. The two types—a tiny bit of moisture and significant leakage—are not concerning. When many potential areas emerge in a single tunnel section, this graded result is very useful. The highest-risk areas will be inspected first by maintenance personnel, who will then choose whether to utilize manual inspection or close-up imaging.

During a long patrol, the inspection crew may compare many tunnel sections with the use of graded output. Although a binary detector may provide a large number of affirmative patches, it cannot identify which parts are more significant. The review queue can be prioritized based on work requirements since the model will establish a normal score and a warning level. For a brief period of time or when there are many moist areas along the same line due to various causes, it is practical. The model will plan a maintenance response based on the context cost and detection confidence.

Interpretability, Robustness and Runtime Analysis

The modules are the subject of an ablation investigation. The deployment-oriented analysis is displayed in Figure 7. The module ablation research is shown in Figure 7(a): eliminating context tokens decreases macro-F1 by 3.8 points; eliminating the Lagrangian penalty raises false alarms by 2.6 points; and eliminating recurrence history reduces confirmed-leakage recall by 2.9 points. This conclusion is supported by uncertainty quantification research [23], as risk-aware decision systems need calibrated confidence under insufficient inspection data.

Robustness to blur, stain reflection, missing context, and invisible tunnel sections is demonstrated in Figure 7(b). Because structural context is still helpful with low-quality pictures, the suggested model shows less deterioration than the visual-only baseline. The stability of the policy update under visual disturbance may also be explained by constrained policy optimization [24]. The whole model takes 21.4 ms per inspection window on an RTX 3060 workstation and 68.7 ms on an industrial edge computer, according to Figure 7(c), which compares runtime performance. Such delay is acceptable for near-real-time tunnel patrol and offline batch inspection, according to Safe Inspection Planning study [25].



Figure 7. Ablation, robustness and runtime analysis. (a) Module ablation. (b) Robustness. (c) Runtime.

Additionally, interpretability analysis reveals that the model takes into account tunnel-zone context, recurrent location history, seam adjacency, and water-boundary continuity. The suggested warning score is evaluated in combination with uncertainty and constraint status since calibration work on neural classifiers [26] shows that probability values might be deceptive if they are not adjusted for the operational shift. By changing the illumination in wet-surface photographs, research on data augmentation [27] enhances resilience. Recent surveys [30], dataset-shift analysis [29], and uncertainty analysis on segmentation [28] have all shown that the inspection model should provide confidence-aware outcomes in many contexts. As a result, the suggested approach serves as a maintenance screening layer by allocating doubtful situations for human inspection, evaluating potential areas, and providing an explanation for warnings.

The ablation findings indicate that the modules serve distinct functions. Context tokens aid in lowering alerts that are structurally unlikely but visually reasonable. A Lagrangian penalty reduces the policy's scope and lowers running expenses. We may identify a slow-onset leak that is not apparent after a single check by looking at recurrence history. Eliminating any one of them will harm another pipeline segment. For leakage detection, the design may thus decide to include perception, context, and limited decision-making. The approach is workable for the real inspection procedure based on the running outcomes. When handled in batches or close to real time, the edge-computer performance is still appropriate for frame-window screening of the patrol video, even if it is slower than workstation inference. Return a prioritized list of places that need manual inspection after scanning a large number of inspection videos in the system. As a result, engineers may concentrate on the crucial places and make fewer observations of numerous areas in the tunnel image. As a result, the model has proven rather helpful so far.

The remaining mistakes mostly occur when the shapes and locations of harmless moisture and actual leakage are similar. There is a water-receding passage at the joint during the maintenance cleaning, and seepage might happen here. A black area next to the cable bracket is thought to indicate a flaw. The circumstances need human verification and special handling. This engineering ambiguity has not been addressed by the model; it has merely documented the reasons for the decision's uncertainty. The warning system may be kept from turning into a black-box label generator by using traceability. Some of the model's inspection modes have been found via robustness testing. Fixed cameras are very useful for recurring history since they have the same region for comparison at various periods. Context tokens and uncertainty are more important because mobile patrol robots have a wider range of viewpoints and lighting conditions. Because manual picture uploads don't include metadata, they must be examined. Because it separates visual evidence and contextually encodes information rather than pushing all inputs into a single feature vector, the framework allows the modes.

A maintenance work order may be created based on the information above. Uncertain instances will be gathered for human review after the patrol, recurrent suspected leaks will be scheduled for follow-up, and high-confidence verified leaks will result in an urgent close-range examination. Compared to a binary leakage label, this progressive process is more useful. Each warning may be linked to a visual area, structural context, uncertainty state, and constraint contribution since it is also accountable. Compared to a purely visual classifier, the features are better suited for engineering applications. The extent of automation may also be determined via deployment analysis. The model will not take the role of structural diagnosis; instead, it suggests an inspection triage function. When determining whether to do maintenance, groundwater conditions, liner deformation, drainage capacity, and repair history should be taken into account even when the algorithm has identified a high-risk person. This flaw is an inevitable security measure for tunnel management, not a flaw in the algorithm. The number of undiscovered flaws and false alarms will decrease with an efficient alarm system, yet the ultimate engineering decision will still be made expertly.

Additionally, this deployment boundary keeps the model from being applied outside of the evidence's purview. Instead of automatically giving a high-confidence label, the system will designate a case for human review if a picture is very grainy, if context fields are missing, or if the tunnel portion has just been repaired. Because the cost of an unjustified automated conclusion is quite large, a cautious approach is appropriate for the functioning of shield tunnels. As a result, the finished product will be used as graded evidence for planning inspections.

Conclusion

A context-aware representation-enhanced PPO-Lagrangian model for shield tunnel leak detection is presented in this research. Create a single model state by combining visual leakage indicators, structural tunnel data, inspection metadata, and recurrence history. Constrained policy optimization is then used to improve detection accuracy and decrease missed-risk instances and false alarms. The findings indicate that the problem of detecting leakage in shield tunnels cannot be viewed as a stand-alone image-labeling task; instead, the meaning of leakage depends on the ring position, joint adjacency, lighting, and past events. The first is to transform a visual recognition problem into a context-aware maintenance decision-making problem. Because the appearance of a leak is often ambiguous and the same mark may have multiple technical meanings at different places or during operation, conversion is necessary. A more useful approach for tunnel inspection systems that need low workload and good sensitivity and reliability has been presented via the integration of feature representation and limited policy learning.

According to the experiment, the suggested model performs better in terms of accuracy, macro-F1 score, confirmed-leakage recall, and warning reliability than both the visual-only baseline and unconstrained reinforcement learning. In the problematic regions of a wet-lined area, seam shadows, reflected water stains, and recurring damp markings, it has a comparatively high number of fringes. The operational restrictions of the policy are satisfied by adding a Lagrangian penalty, and false alarms are decreased by a context-aware representation. This approach is appropriate for batch inspection review and near real-time patrol assistance, according to runtime assessment. In practice, the model is utilized to add an additional step before to physical assessment and after routine picture gathering. Identify situations that need re-inspection, rate suspect locations, filter inspection windows, and provide justifications for warnings. The engineers will be able to concentrate on the sections that provide both structural context and visual indications of danger thanks to this procedure, rather than being excluded at this point. The technique satisfies the criteria for operating urban tunnels since it is fast, has strong accuracy and recall, and has a low false alarm rate.

Whether very low-light or highly obstructed situations need human verification will depend on the quality of the context data. In the future, a comprehensive multi-source leakage detection system will be constructed by combining high-resolution inspection videos, groundwater fluctuation data, tunnel deformation measurements, and maintenance feedback. To identify high-risk areas and provide justifications for sending a leak alert, the model suggested for engineering applications may be used as a traceable decision-support tool. To prevent different operators utilizing different inspection labels, camera paths, and maintenance records, data standardization for all tunnel lines should be introduced in the future. To enhance context encoding and make comparing warning outputs from several modules easier, a single Data Dictionary may be used. Add a few more field tests to confirm that the model remains accurate following maintenance, camera replacement, and seasonal variations in groundwater. The framework has been enhanced by the optimizations, and it can now be linked to other components of a large-scale tunnel health monitoring system, including maintenance planning, structural data, and visual inspection.

Author Contributions

Ursula Stępień contributes to conceptualization, methodology, software, validation, analysis, investigation, data collection, draft preparation, manuscript editing, visualization, supervision, project administration, and funding acquisition. Weronika Błaszczuk contributes to software, validation, analysis. All authors have read and agreed with the manuscript before its submission and publication.

Funding

This research received no specific financial support from any funding agency.

Institutional Review Board Statement

Not applicable.

References

- [1] Huang, H. W., Li, Q. T., & Zhang, D. M. (2018). Deep learning-based image recognition for crack and leakage defects of metro shield tunnel. *Tunnelling and Underground Space Technology*, 77, 166-176. <https://doi.org/10.1016/j.tust.2018.04.002>
- [2] Farrar, C. R., & Worden, K. (2007). An introduction to structural health monitoring. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 365(1851), 303-315. <https://doi.org/10.1098/rsta.2006.1928>
- [3] Cha, Y. J., Choi, W., & Büyüköztürk, O. (2017). Deep learning-based crack damage detection using convolutional neural networks. *Computer-Aided Civil and Infrastructure Engineering*, 32(5), 361-378. <https://doi.org/10.1111/mice.12263>
- [4] Bao, Y., Tang, Z., Li, H., & Zhang, Y. (2019). Computer vision and deep learning-based data anomaly detection method for structural health monitoring. *Structural Health Monitoring*, 18(2), 401-421. <https://doi.org/10.1177/1475921718757405>
- [5] Brownjohn, J. M. W. (2007). Structural health monitoring of civil infrastructure. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 365(1851), 589-622. <https://doi.org/10.1098/rsta.2006.1925>
- [6] Liu, S., Xu, X., Jeon, G., Chen, J., & He, B. G. (2024). Deep learning-based water leakage detection for shield tunnel lining. *Frontiers of Structural and Civil Engineering*, 18(6), 887-898. <https://doi.org/10.1007/s11709-024-1071-5>
- [7] Ronneberger, O., Fischer, P., & Brox, T. (2015). U-Net: Convolutional networks for biomedical image segmentation. *Medical Image Computing and Computer-Assisted Intervention*, 9351, 234-241. https://doi.org/10.1007/978-3-319-24574-4_28
- [8] Chen, L. C., Zhu, Y., Papandreou, G., Schroff, F., & Adam, H. (2018). Encoder-decoder with atrous separable convolution for semantic image segmentation. *European Conference on Computer Vision*, 801-818. https://doi.org/10.1007/978-3-030-01234-2_49
- [9] Liu, Z., Lin, Y., Cao, Y., Hu, H., Wei, Y., Zhang, Z., et al. (2021). Swin Transformer: Hierarchical vision transformer using shifted windows. *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 10012-10022. <https://doi.org/10.1109/ICCV48922.2021.00986>
- [10] Long, J., Shelhamer, E., & Darrell, T. (2015). Fully convolutional networks for semantic segmentation. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 3431-3440. <https://doi.org/10.1109/CVPR.2015.7298965>
- [11] Mnih, V., Badia, A. P., Mirza, M., Graves, A., Lillicrap, T., Harley, T., et al. (2016). Asynchronous methods for deep reinforcement learning. *International Conference on Machine Learning*, 48, 1928-1937. <https://doi.org/10.5555/3045390.3045594>
- [12] Schulman, J., Levine, S., Abbeel, P., Jordan, M., & Moritz, P. (2015). Trust region policy optimization. *International Conference on Machine Learning*, 37, 1889-1897. <https://doi.org/10.5555/3045118.3045319>
- [13] Achiam, J., Held, D., Tamar, A., & Abbeel, P. (2017). Constrained policy optimization. *International Conference on Machine Learning*, 70, 22-31. <https://doi.org/10.5555/3305381.3305384>
- [14] Brunke, L., Greeff, M., Hall, A. W., Yuan, Z., Zhou, S., Panerati, J., & Schoellig, A. P. (2022). Safe learning in robotics: From learning-based control to safe reinforcement learning. *Annual Review of Control, Robotics, and Autonomous Systems*, 5, 411-444. <https://doi.org/10.1146/annurev-control-042920-020211>
- [15] Wang, W., Su, C., Han, G., Dong, Y., & Li, Y. (2024). Efficient segmentation of water leakage in shield tunnel lining with convolutional neural network. *Structural Health Monitoring*, 23(2), 671-685. <https://doi.org/10.1177/14759217231171696>
- [16] He, K., Zhang, X., Ren, S., & Sun, J. (2016). Deep residual learning for image recognition. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 770-778. <https://doi.org/10.1109/CVPR.2016.90>
- [17] Song, J., He, L. X., & Long, H. P. (2023). Tunnel lining multi-defect detection based on an improved You Only Look Once Version 7 algorithm. *IEEE Access*, 11, 125171-125184. <https://doi.org/10.1109/ACCESS.2023.3330843>
- [18] Lin, T. Y., Goyal, P., Girshick, R., He, K., & Dollár, P. (2017). Focal loss for dense object detection. *Proceedings of the IEEE International Conference on Computer Vision*, 2980-2988. <https://doi.org/10.1109/ICCV.2017.324>
- [19] Cui, Y., Jia, M., Lin, T. Y., Song, Y., & Belongie, S. (2019). Class-balanced loss based on effective number of samples. *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 9268-9277. <https://doi.org/10.1109/CVPR.2019.00949>

- [20] Berman, M., Triki, A. R., & Blaschko, M. B. (2018). The Lovász-Softmax loss: A tractable surrogate for the optimization of the intersection-over-union measure in neural networks. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 4413-4421. <https://doi.org/10.1109/CVPR.2018.00464>
- [21] Gneiting, T., & Raftery, A. E. (2007). Strictly proper scoring rules, prediction, and estimation. *Journal of the American Statistical Association*, 102(477), 359-378. <https://doi.org/10.1198/016214506000001437>
- [22] Van Calster, B., McLernon, D. J., van Smeden, M., Wynants, L., & Steyerberg, E. W. (2019). Calibration: The Achilles heel of predictive analytics. *BMC Medicine*, 17, 230. <https://doi.org/10.1186/s12916-019-1466-7>
- [23] Abdar, M., Pourpanah, F., Hussain, S., Rezazadegan, D., Liu, L., Ghavamzadeh, M., et al. (2021). A review of uncertainty quantification in deep learning: Techniques, applications and challenges. *Information Fusion*, 76, 243-297. <https://doi.org/10.1016/j.inffus.2021.05.008>
- [24] Pakdaman Naeini, M., Cooper, G., & Hauskrecht, M. (2015). Obtaining well calibrated probabilities using Bayesian binning. *Proceedings of the AAAI Conference on Artificial Intelligence*, 29(1), 2901-2907. <https://doi.org/10.1609/aaai.v29i1.9602>
- [25] García, J., & Fernández, F. (2015). A comprehensive survey on safe reinforcement learning. *Journal of Machine Learning Research*, 16, 1437-1480. <https://doi.org/10.5555/2789272.2886795>
- [26] Selvaraju, R. R., Cogswell, M., Das, A., Vedantam, R., Parikh, D., & Batra, D. (2017). Grad-CAM: Visual explanations from deep networks via gradient-based localization. *Proceedings of the IEEE International Conference on Computer Vision*, 618-626. <https://doi.org/10.1109/ICCV.2017.74>
- [27] Shorten, C., & Khoshgoftaar, T. M. (2019). A survey on image data augmentation for deep learning. *Journal of Big Data*, 6, 60. <https://doi.org/10.1186/s40537-019-0197-0>
- [28] Nair, T., Precup, D., Arnold, D. L., & Arbel, T. (2020). Exploring uncertainty measures in deep networks for multiple sclerosis lesion detection and segmentation. *Medical Image Analysis*, 59, 101557. <https://doi.org/10.1016/j.media.2019.101557>
- [29] Kuleshov, V., & Ermon, S. (2017). Estimating uncertainty online against an adversary. *Proceedings of the AAAI Conference on Artificial Intelligence*, 31(1), 2110-2116. <https://doi.org/10.1609/aaai.v31i1.10949>
- [30] Gawlikowski, J., Tassi, C. R. N., Ali, M., Lee, J., Humt, M., Feng, J., et al. (2023). A survey of uncertainty in deep neural networks. *Artificial Intelligence Review*, 56, 1513-1589. <https://doi.org/10.1007/s10462-023-10562-9>