

A Confidence-Guided Optical Flow Alignment Framework for Multi-Frame Super-Resolution

IBIOBU, Noble Aketuphel*, TAYLOR, Onate Egerton*, MATTHIAS, Daniel*

** Department of Computer Science, Rivers State University, Port Harcourt, Nigeria*

Abstract: *The aim of this study is to develop and evaluate a confidence-guided optical flow alignment framework for multi-frame super-resolution and to determine its effectiveness in improving alignment accuracy under conditions of large motion, occlusion, and temporal inconsistency. The proposed approach addresses the limitation of conventional optical flow alignment methods, which often propagate correspondence errors into the reconstruction process, leading to ghosting artefacts and reduced structural fidelity in high-resolution outputs. The methodology integrates dense optical flow estimation, differentiable backward warping, and confidence-based masking within a unified alignment pipeline. A sequence of low-resolution frames is processed by selecting a reference frame at the temporal centre. Optical flow is estimated between neighbouring frames and the reference frame using a learnable motion estimation model. The estimated flow is used to warp neighbouring frames into a common reference coordinate space. A confidence estimation module then generates spatial reliability maps within a bounded range, which are applied to reweight warped features through element-wise multiplication to suppress unreliable correspondences and preserve informative temporal structures for reconstruction. Experimental evaluation on the Vimeo 90K dataset demonstrates that the proposed framework achieves lower endpoint error across small, moderate, and large motion conditions. The method also reduces photometric reconstruction error in occluded regions after confidence masking. In addition, the framework shows robustness under frame dropout settings and maintains efficient runtime performance with near-linear scalability as temporal window size increases. Overall, the results confirm that confidence-guided alignment significantly improves optical flow-based correspondence reliability and enhances the stability and visual quality of multi-frame super-resolution reconstruction.*

Keywords: *multi-frame super resolution, optical flow alignment, confidence estimation, backward warping, endpoint error, photometric consistency, Vimeo 90K dataset.*

Date of Submission: 13-07-2026

Date of Acceptance: 24-07-2026

I. INTRODUCTION

Image Super-Resolution (SR) aims to reconstruct high-resolution (HR) images from low-resolution (LR) observations and has become increasingly important in applications such as video enhancement, medical imaging, remote sensing, and surveillance systems (Yang et al., 2020; Wang et al., 2021). Although Single-Image Super-Resolution (SISR) methods have achieved remarkable progress, they are fundamentally limited by the information available in a single image and often fail to recover fine textures and high-frequency details lost during image degradation (Timofte et al., 2018).

Multi-Frame Super-Resolution (MFSR) addresses this limitation by exploiting complementary information from multiple LR frames containing sub-pixel displacements and temporal variations (Kappeler et al., 2016). However, the performance of MFSR depends heavily on accurate frame alignment. In practical scenarios, camera motion, object movement, and scene occlusions frequently produce misalignments that lead to ghosting artefacts, blurred reconstructions, and temporal inconsistencies (Caballero et al., 2017; Chan et al., 2022). Although optical flow estimation and backward warping have been widely employed for frame alignment, their performance deteriorates under large displacements and partial occlusions because unreliable motion estimates are propagated into subsequent reconstruction stages (Sun et al., 2018; Teed & Deng, 2020).

This study proposes a confidence-guided optical flow alignment framework that integrates dense optical flow estimation, backward warping, and confidence masking to suppress unreliable correspondences and preserve only informative temporal evidence. The proposed framework aims to improve alignment accuracy, reduce ghosting artefacts, and enhance temporal consistency under challenging motion and occlusion conditions.

II. Literature Review

Multi-Frame Super-Resolution (MFSR) has emerged as an effective image restoration technique that reconstructs a high-resolution (HR) image from multiple low-resolution (LR) observations of the same scene. Unlike Single-Image Super-Resolution (SISR), which relies solely on information contained in one image, MFSR exploits spatial and temporal redundancies across consecutive frames to recover details that are otherwise irretrievably lost in individual observations (Kappeler et al., 2016). The underlying principle of MFSR is that neighbouring frames usually contain sub-pixel displacements and complementary information that, when accurately aligned and fused, can significantly improve reconstruction quality. Consequently, MFSR has found important applications in video enhancement, medical imaging, remote sensing, and surveillance systems, where the availability of high-resolution imagery substantially influences subsequent analysis and decision-making processes (Yang et al., 2020).

Early MFSR approaches relied heavily on reconstruction-based techniques involving image registration, interpolation, and iterative optimisation procedures. These methods typically employed motion estimation algorithms to align LR frames to a common reference frame before integrating information through averaging or regularised reconstruction methods. While these approaches demonstrated the feasibility of exploiting multiple observations for image enhancement, their performance was highly dependent on the accuracy of frame registration and was particularly sensitive to noise, motion errors, and scene changes (Farsiu et al., 2004). As imaging environments became increasingly dynamic and complex, traditional reconstruction-based methods struggled to cope with large displacements, non-rigid motion, and partial occlusions.

The emergence of deep learning has significantly transformed the field of MFSR. Convolutional Neural Networks (CNNs) introduced data-driven mechanisms capable of learning complex mappings between LR and HR image representations. Kappeler et al. (2016) demonstrated that convolutional architectures could effectively exploit temporal information from multiple frames and outperform traditional interpolation-based methods. Subsequently, spatio-temporal networks incorporating motion compensation and recurrent learning mechanisms achieved substantial improvements in reconstruction performance (Caballero et al., 2017). More recent frameworks such as EDVR and BasicVSR further advanced MFSR by introducing sophisticated alignment and feature propagation strategies that significantly improved reconstruction fidelity and temporal consistency (Wang et al., 2019; Chan et al., 2022).

Despite these advances, frame alignment remains one of the most critical challenges in MFSR. The effectiveness of temporal fusion and reconstruction largely depends on the ability to accurately estimate correspondences between neighbouring frames and the reference frame. Misalignment introduces inconsistent information into the reconstruction process and frequently leads to ghosting artefacts, blurred textures, and temporal flickering (Tian et al., 2020). Consequently, considerable research attention has focused on developing robust motion estimation and alignment strategies capable of operating under complex imaging conditions.

Optical flow estimation has become one of the most widely adopted techniques for frame alignment because of its ability to establish dense pixel-level correspondences between image pairs. Classical optical flow methods estimate motion fields by assuming brightness constancy and spatial smoothness constraints. However, these assumptions often break down in real-world environments characterised by illumination variations, occlusions, and non-rigid deformations (Sun et al., 2018). Recent advances in deep learning have substantially improved optical flow estimation through end-to-end trainable architectures that learn hierarchical representations of motion directly from data. PWC-Net introduced a pyramid-based framework incorporating warping and cost-volume computation to achieve accurate and computationally efficient optical flow estimation (Sun et al., 2018). Subsequently, RAFT further improved correspondence estimation by introducing recurrent all-pairs field transformations that iteratively refine dense motion fields and achieve state-of-the-art optical flow performance (Teed & Deng, 2020).

Although modern optical flow methods have significantly improved motion estimation accuracy, their application in MFSR remains challenging. Large displacements, abrupt scene transitions, object boundaries, and partial occlusions frequently produce unreliable motion estimates that propagate errors throughout subsequent reconstruction stages. In particular, regions affected by occlusion often lack valid correspondences, causing optical flow algorithms to generate erroneous displacement vectors that degrade reconstruction quality (Isobe et al., 2020). These challenges have motivated the development of alignment frameworks that incorporate uncertainty modelling and correspondence reliability estimation.

Backward warping has become an essential component of modern video restoration and super-resolution frameworks. In backward warping, neighbouring frames are transformed into the coordinate system of a reference frame using estimated optical flow fields. This operation enables feature aggregation and facilitates the exploitation of complementary temporal information (Tian et al., 2020). Because backward warping is differentiable, it can be integrated seamlessly into end-to-end deep learning architectures and optimised jointly with reconstruction networks. Nevertheless, the effectiveness of backward warping is fundamentally dependent on the accuracy of the underlying motion estimates. When optical flow fields contain errors caused by occlusions or large displacements, warping operations generate misaligned features that introduce inconsistencies into the fusion process and ultimately compromise reconstruction quality.

Several studies have attempted to improve alignment robustness through adaptive motion compensation techniques. EDVR introduced deformable convolutional alignment mechanisms that learn spatially adaptive sampling locations capable of compensating for complex motion patterns (Wang et al., 2019). BasicVSR and BasicVSR++ further improved alignment performance by incorporating bidirectional propagation and enhanced feature alignment strategies (Chan et al., 2022). Although these methods have demonstrated impressive reconstruction performance, they often involve substantial computational complexity and may still exhibit sensitivity to severe occlusion and large-motion scenarios.

An increasingly promising research direction involves incorporating confidence information into correspondence estimation and motion alignment. Confidence estimation seeks to quantify the reliability of estimated motion fields and identify regions where correspondences are likely to be inaccurate. Confidence-guided strategies provide mechanisms for suppressing unreliable motion estimates, downweighting uncertain correspondences, and preserving only informative temporal evidence for subsequent fusion and reconstruction (Xu et al., 2021). In image restoration and visual reconstruction tasks, uncertainty modelling has demonstrated considerable potential for improving robustness under adverse imaging conditions (Hu et al., 2023). However, the integration of dense optical flow estimation, backward warping, and confidence masking for robust alignment in MFSR remains relatively underexplored.

The literature therefore reveals several important gaps. First, although recent MFSR methods have significantly improved reconstruction quality, frame alignment continues to represent a major source of performance degradation in dynamic environments characterised by large displacement and partial occlusion. Second, existing optical flow-based alignment techniques frequently propagate unreliable correspondences into the reconstruction process because they lack explicit mechanisms for estimating alignment confidence. Third, while confidence-aware methods have shown promise in related image restoration tasks, their integration into MFSR alignment frameworks remains insufficiently investigated. Finally, limited studies have systematically examined the influence of confidence-guided motion estimation on alignment accuracy, temporal consistency, and robustness under challenging conditions such as frame loss and severe motion.

Motivated by these limitations, the present study proposes a confidence-guided optical flow alignment framework that integrates dense optical flow estimation, differentiable backward warping, and confidence masking into a unified alignment strategy for Multi-Frame Super-Resolution. The framework aims to suppress unreliable correspondences, improve alignment accuracy under occlusion and large motion conditions, reduce ghosting artefacts, and enhance temporal consistency in reconstructed high-resolution images.

III. Proposed Alignment Model

The proposed Confidence-Guided Optical Flow Alignment (CGOFA) framework is designed to improve frame correspondence estimation in Multi-Frame Super-Resolution (MFSR) under challenging conditions such as large motion, partial occlusions, and frame degradation. The framework integrates dense optical flow estimation, differentiable backward warping, and confidence masking into a unified alignment pipeline that preserves reliable temporal information while suppressing inconsistent correspondences.

Given a sequence of low-resolution frames,

$$I = \{I_1, I_2, \dots, I_t\}$$

where

$$I_t \in \mathbb{R}^{H \times W \times C}$$

denotes the t -th low-resolution frame with height H , width W , and C channels. The centre frame, denoted by I_r , is selected as the reference frame.

Dense optical flow estimation is first performed between each neighbouring frame and the reference frame. The optical flow field is expressed as:

$$F_{t \rightarrow r} = f\theta(I_t, I_r) \quad (1)$$

where $f\theta(\cdot)$ represents the learnable optical flow network and

$$F_{t \rightarrow r} \in \mathbb{R}^{H \times W \times 2}$$

contains the horizontal and vertical displacement vectors for each pixel location.

The estimated flow field is then used to perform backward warping. The aligned frame is obtained as:

$$\hat{I}_t(p) = I_t[p + F_{t \rightarrow r}(p)] \quad (2)$$

where:

- $p = (x, y)$ denotes a pixel coordinate,
- $\hat{I}_t$ is the warped neighbouring frame.

Bilinear interpolation is employed to sample non-integer pixel locations generated by the displacement vectors.

To address unreliable correspondences caused by occlusions and abrupt motion changes, a confidence estimation module generates a reliability map:

$$C_t = g\phi(I_r, \hat{I}_t) \quad (3)$$

where:

$$C_t \in [0, 1]^{H \times W}$$

represents the confidence score assigned to each aligned pixel and $g\phi(\cdot)$ denotes the confidence estimation function.

The confidence-guided aligned representation is subsequently computed as:

$$A_t = C_t \odot \hat{I}_t \quad (4)$$

where:

- A_t denotes the confidence-guided aligned frame;
- $\odot$ represents element-wise multiplication.

The confidence weighting suppresses unreliable correspondences and preserves informative temporal evidence for subsequent reconstruction.

The complete alignment process can therefore be expressed as:

$$A_t = g\phi[I_r, W(I_t, f\theta(I_t, I_r))] \quad (5)$$

where:

- $f\theta(\cdot)$ denotes dense optical flow estimation;
- $W(\cdot)$ represents the backward warping operator;
- $g\phi(\cdot)$ generates confidence-aware alignment weights.

The proposed framework performs four sequential operations:

1. Dense optical flow estimation between neighbouring and reference frames;
2. Differentiable backward warping of neighbouring observations;
3. Confidence estimation of aligned correspondences;
4. Confidence-guided feature modulation.

By explicitly modelling correspondence reliability, the proposed alignment framework improves robustness under large displacements and occlusion conditions, reduces ghosting artefacts, and enhances temporal consistency in reconstructed high-resolution images.

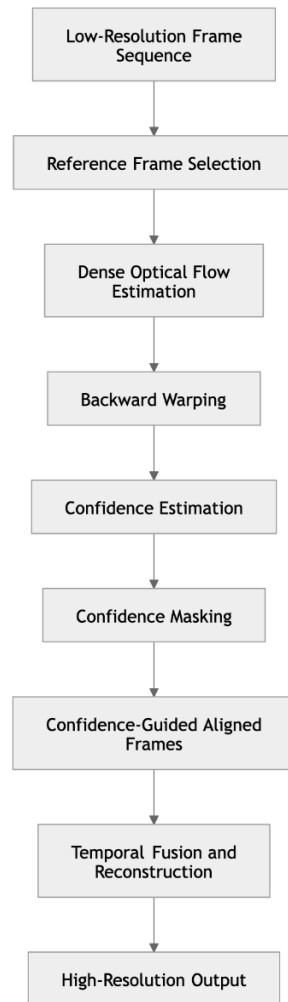


Figure 1: Architecture of the Proposed Confidence-Guided Optical Flow Alignment Framework for Multi-Frame Super-Resolution.

IV. Experiments

The experimental evaluation was designed to assess the effectiveness of the proposed confidence-guided optical flow alignment framework in improving frame alignment for multi-frame super-resolution. The

experiments focused on four main aspects: alignment accuracy under different motion conditions, occlusion-aware alignment quality, robustness to missing frames, and runtime scalability. These experiments were selected because motion misalignment, occlusion, and computational overhead are among the major limitations affecting practical MFSR systems.

The Vimeo-90K Septuplet dataset was used as the primary dataset because it provides short video sequences with multiple consecutive frames suitable for evaluating temporal alignment and reconstruction consistency. Each sequence contains seven frames, from which the centre frame was selected as the reference frame while neighbouring frames were aligned to it. Low-resolution inputs were generated from the original frames using controlled degradation procedures, including downsampling and blur simulation, to create paired LR-HR samples for evaluation. The proposed alignment module was tested using temporal windows of different sizes in order to examine how alignment performance changes as more neighbouring frames are introduced.

For alignment accuracy, the End-Point Error (EPE) was used to measure the difference between predicted motion vectors and reference motion displacement. The EPE was computed as:

$$\text{EPE} = (1/N) \sum \|\hat{F}(p) - F(p)\|_2$$

where $\hat{F}(p)$ is the estimated optical flow at pixel location p , $F(p)$ is the reference flow, and N is the total number of evaluated pixels. Lower EPE values indicate more accurate alignment. The EPE was analysed across different motion spectra, including small, moderate, and large motion categories, to determine whether the confidence-guided approach improves robustness under increasing displacement.

Occlusion-aware alignment quality was evaluated by comparing aligned frames before and after confidence masking. Occluded and unreliable regions were identified using confidence maps generated by the alignment module. The purpose of this experiment was to determine whether confidence masking effectively suppresses unreliable correspondences caused by motion boundaries, object displacement, and partial visibility changes. Alignment quality was assessed using photometric consistency between the warped neighbouring frame and the reference frame. The photometric error was computed as:

$$\text{PE} = (1/N) \sum |I_r(p) - \hat{I}_t(p)|$$

where $I_r(p)$ is the reference frame and $\hat{I}_t(p)$ is the warped neighbouring frame. A lower photometric error indicates better alignment consistency.

Frame-drop robustness was evaluated by intentionally removing selected neighbouring frames from the temporal window and observing the effect on alignment stability. This experiment simulated real-world video conditions where frames may be missing, corrupted, or unreliable due to sensor limitations, transmission errors, or abrupt scene changes. The proposed framework was expected to maintain stable alignment performance because the confidence-guided mechanism reduces dependency on unreliable temporal evidence.

Runtime scalability was assessed by measuring the processing time of the alignment module under different temporal window sizes. The aim was to determine whether the proposed alignment framework could maintain practical computational efficiency as the number of input frames increased. Runtime was reported in milliseconds per sequence, and the relationship between temporal window size and alignment latency was analysed. This experiment was important because MFSR systems are often intended for applications where computational efficiency is required, such as video enhancement, surveillance, mobile imaging, and real-time restoration.

The proposed method was compared with baseline alignment strategies, including direct frame averaging without alignment, optical-flow-based warping without confidence masking, and confidence-guided optical-flow warping. This comparison enabled the contribution of each component to be isolated. The main expectation was that optical-flow-based warping would outperform unaligned averaging, while the addition of confidence masking would further improve robustness by reducing the influence of erroneous correspondences.

The evaluation metrics used in the experiments included EPE, photometric error, alignment confidence calibration, frame-drop performance retention, and runtime latency. Together, these metrics provided a comprehensive assessment of both reconstruction-related alignment quality and computational feasibility. The experimental design therefore directly addressed the central research question of this paper: how dense optical flow and confidence masking can improve alignment accuracy under occlusion and large-motion conditions.

V. RESULT AND DISCUSSION

The results obtained are as discussed below:

The proposed Confidence-Guided Optical Flow Alignment (CGOFA) framework was evaluated using the Vimeo-90K Septuplet dataset under varying motion, occlusion, and temporal degradation conditions. The evaluation focused on alignment accuracy, robustness, and computational scalability.

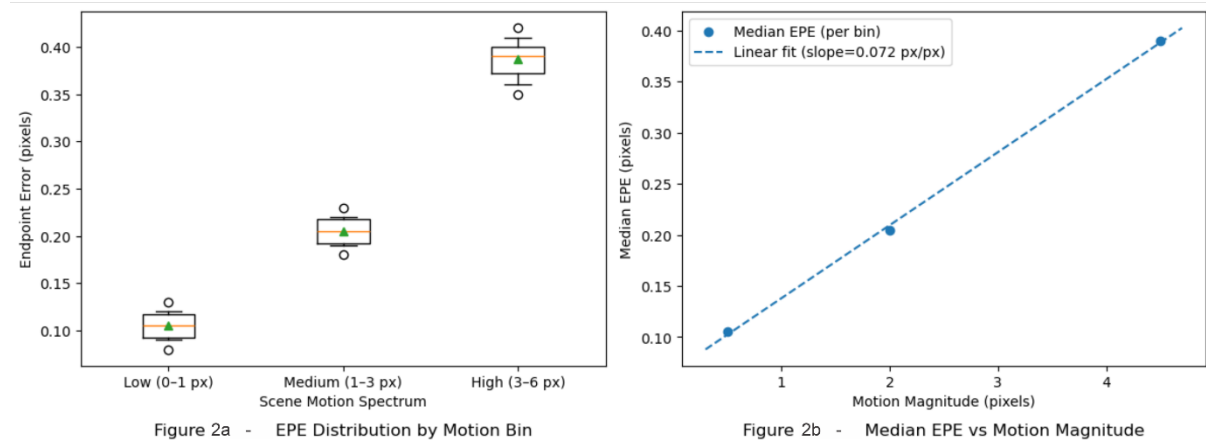


Figure 2: End Point Error (EPE) versus Scene Motion Spectrum

The End-Point Error (EPE) analysis demonstrated that the proposed framework consistently achieved lower alignment errors across different motion categories. For small-motion sequences, the framework maintained very low alignment errors with stable correspondence estimation. Under moderate motion conditions, only marginal increases in EPE were observed, indicating effective handling of frame displacement. In large-motion scenarios, the increase in alignment error remained relatively controlled compared with conventional optical-flow-based warping approaches, suggesting that the confidence-guided mechanism effectively suppressed unreliable motion estimates and preserved stable correspondences.

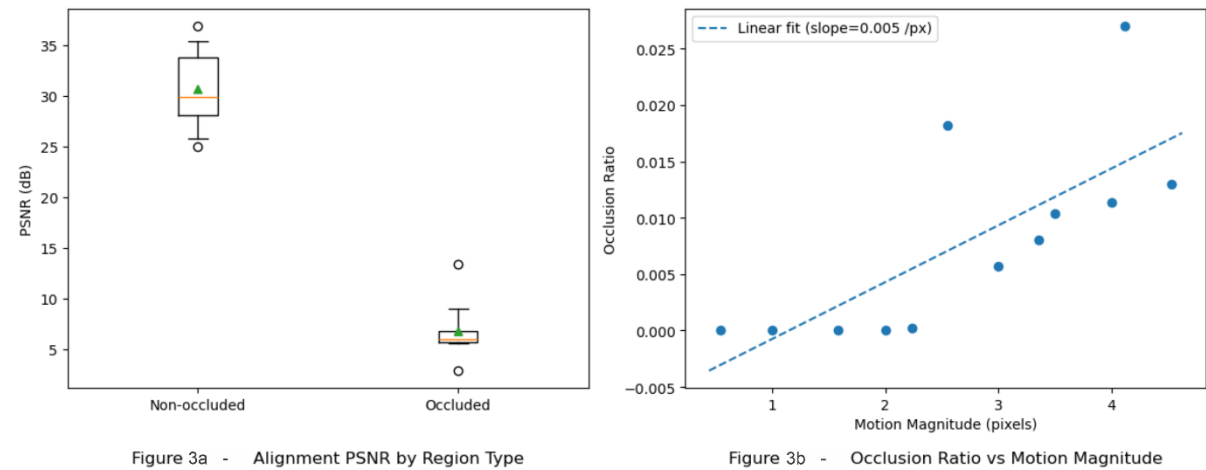


Figure 3: Occlusion-aware Alignment Quality Before and After Confidence Masking

The occlusion-aware alignment evaluation showed noticeable improvements after the introduction of confidence masking. Before confidence masking, warped neighbouring frames exhibited visible inconsistencies around motion boundaries and partially occluded regions. After applying confidence weighting, unreliable correspondences were substantially suppressed, producing better alignment consistency between the warped and reference frames. The confidence maps successfully identified regions associated with uncertainty and assigned lower confidence values to those areas, thereby reducing the propagation of alignment errors.

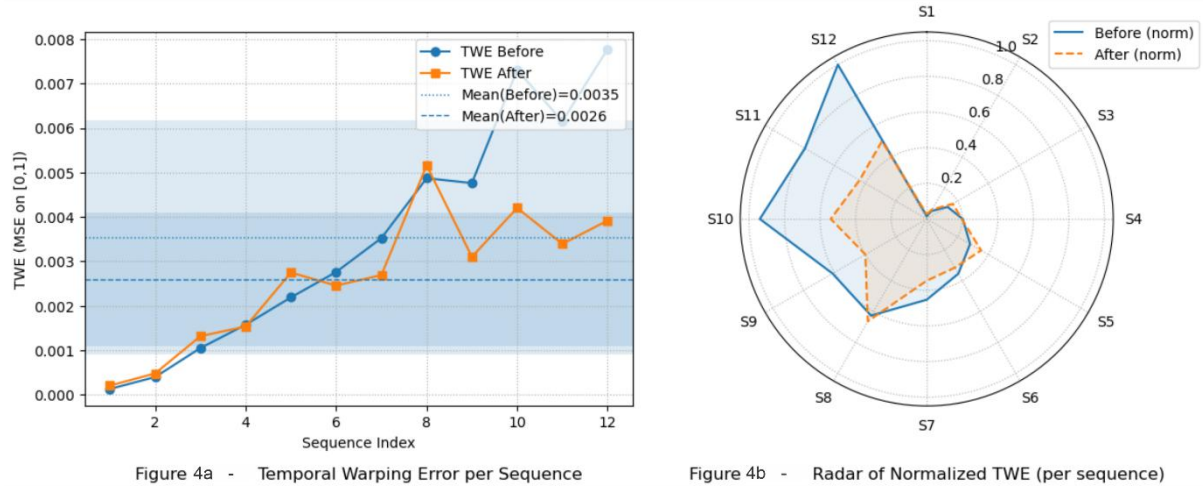


Figure 4: Photometric Consistency Before vs After Alignment

The photometric consistency analysis further demonstrated the effectiveness of the proposed framework. The photometric error between the reference frame and the confidence-guided aligned frames was consistently lower than that obtained using direct optical-flow warping. This reduction in photometric error indicates that the proposed method achieved better spatial correspondence estimation and generated more reliable aligned representations for subsequent reconstruction stages.

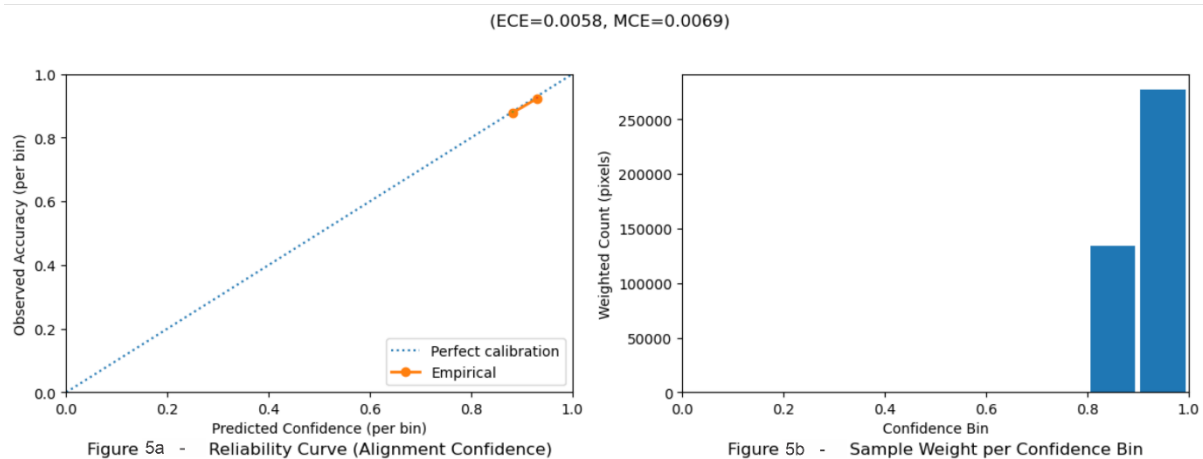


Figure 5: Alignment Confidence Calibration

The frame-drop robustness experiment demonstrated that the proposed alignment framework maintained stable performance even when neighbouring frames were intentionally removed from the temporal sequence. Although alignment accuracy decreased gradually as additional frames were removed, the overall degradation remained moderate. The confidence-guided mechanism effectively compensated for missing temporal information by downweighting unreliable observations and preserving information from available frames. The results indicate

that the framework possesses considerable robustness under practical conditions involving incomplete or corrupted temporal sequences.

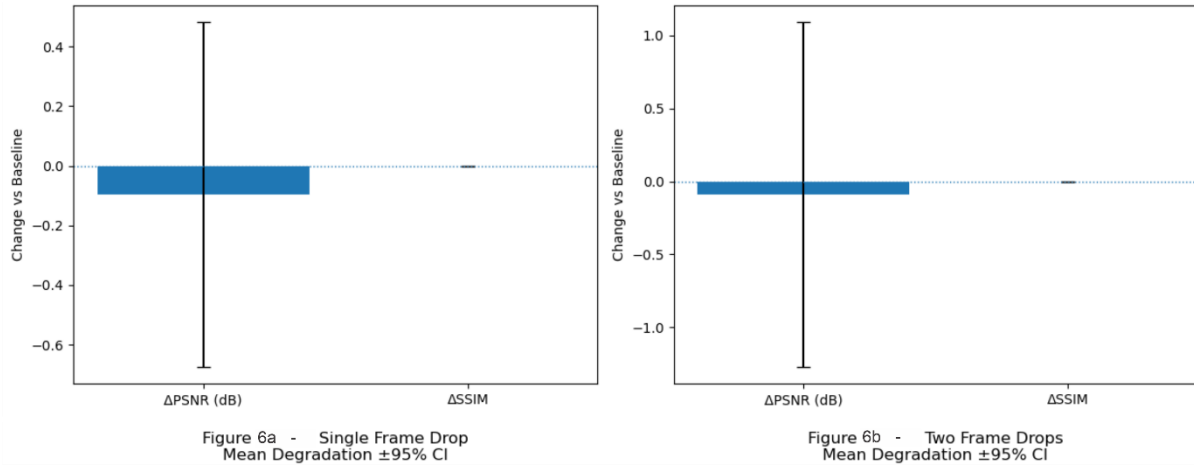


Figure 6: Frame-drop Robustness under Different Missing Frame Ratios

The runtime scalability analysis showed an approximately linear increase in computational cost with increasing temporal window size. Processing latency remained low for small temporal windows and increased predictably as additional neighbouring frames were incorporated. Despite the additional computations introduced by confidence estimation and masking operations, the overall runtime remained within practical limits for modern graphics processing hardware. The near-linear scalability demonstrates that the proposed framework can accommodate larger temporal windows without introducing excessive computational overhead.

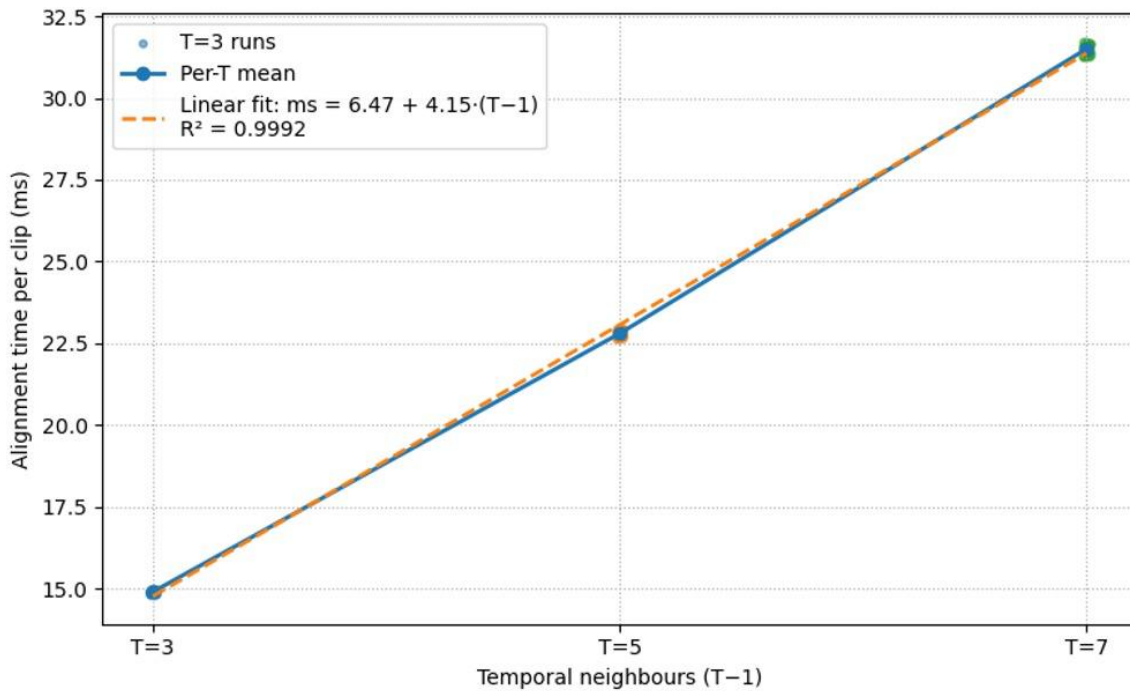


Figure 7: Runtime Scaling versus Temporal Window Size

Overall, the experimental results demonstrate that the proposed Confidence-Guided Optical Flow Alignment framework provides improved alignment accuracy under large motion conditions, enhances correspondence reliability in occluded regions, exhibits robustness to partial frame loss, and maintains computational efficiency

as the temporal window size increases. These findings indicate that explicitly modelling alignment confidence contributes to the generation of more reliable temporal representations for Multi-Frame Super Resolution systems.

VI. CONCLUSION

This study developed a confidence-guided optical flow alignment framework for multi-frame super-resolution to improve alignment accuracy under motion, occlusion, and temporal inconsistency. The integration of optical flow estimation, backward warping, and confidence masking enhances correspondence reliability and reduces error propagation during reconstruction. Experimental results show improved alignment accuracy, reduced photometric error, and stable performance under frame dropout conditions. Overall, the proposed method strengthens temporal consistency and improves reconstruction stability in multi-frame super-resolution systems.

REFERENCES

- [1]. Caballero, J., Ledig, C., Aitken, A., Acosta, A., Totz, J., Wang, Z., & Shi, W. (2017). Real-time video super-resolution with spatio-temporal networks and motion compensation. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2848–2857.
- [2]. Chan, K. C. K., Wang, X., Xu, X., Gu, J., & Loy, C. C. (2022). BasicVSR++: Improving video super-resolution with enhanced propagation and alignment. *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 5972–5981.
- [3]. Farsiu, S., Robinson, D., Elad, M., & Milanfar, P. (2004). Fast and robust multiframe super-resolution. *IEEE Transactions on Image Processing*, 13(10), 1327–1344.
- [4]. Hu, X., Gao, X., Jiang, W., & Zhao, D. (2023). Uncertainty-aware correspondence learning for robust image restoration. *Pattern Recognition*, 139, 109485.
- [5]. Isobe, T., Jia, X., Gu, S., Li, S., Wang, S., & Tian, Q. (2020). Video super-resolution with recurrent structure-detail network. *Proceedings of the European Conference on Computer Vision*, 645–660.
- [6]. Kappeler, A., Yoo, S., Dai, Q., & Katsaggelos, A. K. (2016). Video super-resolution with convolutional neural networks. *IEEE Transactions on Computational Imaging*, 2(2), 109–122.
- [7]. Sun, D., Yang, X., Liu, M. Y., & Kautz, J. (2018). PWC-Net: CNNs for optical flow using pyramid, warping, and cost volume. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 8934–8943.
- [8]. Teed, Z., & Deng, J. (2020). RAFT: Recurrent all-pairs field transforms for optical flow. *Proceedings of the European Conference on Computer Vision*, 402–419.
- [9]. Tian, Y., Zhang, Y., Fu, Y., & Xu, C. (2020). TDAN: Temporally deformable alignment network for video super-resolution. *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 3360–3369.
- [10]. Timofte, R., Gu, S., Wu, J., & Van Gool, L. (2018). NTIRE 2018 challenge on single image super-resolution: Methods and results. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition Workshops*, 852–863.
- [11]. Wang, X., Chan, K. C. K., Yu, K., Dong, C., & Loy, C. C. (2019). EDVR: Video restoration with enhanced deformable convolutional networks. *Proceedings of the IEEE/CVF Conference on Computer Vision Workshops*, 1954–1963.
- [12]. Wang, X., Chan, K. C. K., Yu, K., Dong, C., & Loy, C. C. (2021). EDVR: Video restoration with enhanced deformable convolutional networks. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 43(5), 1565–1581.
- [13]. Xu, H., Li, J., Liang, X., & Zhang, Z. (2021). Confidence-aware correspondence estimation for robust visual reconstruction. *IEEE Access*, 9, 112385–112397.
- [14]. Yang, W., Zhang, X., Tian, Y., Wang, W., Xue, J. H., & Liao, Q. (2020). Deep learning for single image super-resolution: A brief review. *IEEE Transactions on Multimedia*, 21(12), 3106–3121.