



This MICCAI paper is the Open Access version, provided by the MICCAI Society. It is identical to the accepted version, except for the format and this watermark; the final published version is available on SpringerLink.

Domain-Adapted Bi-Planar RAFT for Dense 3D Cardiac Motion Estimation

Nubia Correa¹, Ryan Tan², Sarah J. Zou¹, Garry Chinn¹, and Craig S. Levin¹

¹ Molecular Imaging Instrumentation Laboratory, Department of Electrical Engineering, Stanford University, Stanford, CA, USA
necorrea@stanford.edu, sjzou@stanford.edu, gchinn@stanford.edu, cslevin@stanford.edu

² Department of Computer Science, Stanford University, Stanford, CA, USA
tanryan@stanford.edu

Abstract. Cardiac PET imaging suffers from motion artifacts due to long acquisition times and physiological motion. As a proof of concept, we investigate dense 3D cardiac motion estimation by bi-planar fusion of two orthogonally fine-tuned 2D RAFT (Recurrent All-Pairs Field Transforms) models on a simulated 4D XCAT phantom. We fine-tune a coronal-view model (Y-view) to predict (dx, dz) and a sagittal-view model (X-view) to predict (dy, dz) , then deterministically fuse their outputs into a volumetric flow field (dx, dy, dz) . On the held-out evaluation pair from the phantom, the Y- and X-view models achieve 0.148 and 0.204 px EPE, respectively, while the fused field achieves 0.178 px EPE (0.27 mm), 96.7% sub-pixel accuracy, and 88.6% directional accuracy. The two dz estimates are correlated ($r=0.880$) and have similar signed bias (-0.0091 and -0.0085 voxels), and averaging reduces dz MAE to 0.0504 voxels. Relative to the evaluated Farneback and unadapted pretrained-RAFT baselines, fusion reduces 3D EPE by 74% and 66%, respectively. These results support the within-phantom feasibility of parameter-free bi-planar fusion of two 2D RAFT models, while generalization to independent anatomies, realistic PET statistics, and clinical data remains to be established.

Keywords: Cardiac PET · Optical Flow · Domain Adaptation · Motion Correction · XCAT Phantom

1 Introduction

Cardiac PET enables quantitative assessment of myocardial perfusion and metabolism, but physiological motion degrades spatial resolution and quantitative accuracy during multi-minute acquisitions. Gated reconstruction reduces these effects by dividing data into cardiac phases, but accurate inter-gate registration remains important for static reconstruction and motion analysis [7].

RAFT [11] provides accurate dense 2D optical flow through all-pairs correlation and recurrent refinement, but models trained on natural images face substantial domain differences in medical imaging. Volumetric registration can

instead use dedicated 3D architectures [1], while 2.5D, multi-view, and biplanar approaches exploit complementary planes [13, 12, 3].

Motivated by this prior work, we evaluate a cardiac-PET-specific proof of concept in which orthogonally fine-tuned 2D RAFT models provide complementary motion components that are fused into a 3D field. The present study uses one noise-free XCAT anatomy [9] with exact simulated motion.

Contributions. (i) A bi-planar fusion strategy combining complementary components from orthogonal 2D RAFT models; (ii) a within-phantom slice-count analysis showing diminishing returns beyond 16 slices per orientation; (iii) analysis of the shared dz component using correlation, bias, and error; and (iv) comparison with Farneback and unadapted pretrained RAFT as reference baselines.

2 Related Work

Optical flow estimation has a long history in medical image analysis, spanning classical variational methods to modern deep learning approaches. Early cardiac motion estimation relied on Horn-Schunck [6] and Lucas-Kanade formulations, extended with mass conservation constraints for PET [2] and elastic registration for multi-gate imaging [5]. More recent approaches combine MR-based motion tracking with PET reconstruction [8] and joint estimation with denoising of parametric images [14]. Deep learning advanced optical flow through FlowNet [4], PWC-Net [10], and RAFT [11], while learning-based volumetric registration methods such as VoxelMorph [1] directly estimate 3D deformation fields. Between purely 2D and full 3D processing, 2.5D methods combine information from multiple slices or views [13]. Related registration work has also exploited biplanar radiographs [12] and two orthogonal projections [3] to constrain 3D motion or deformation. Our contribution is therefore narrower than the general use of orthogonal views: we evaluate whether independently fine-tuned 2D RAFT models can provide complementary displacement components for bi-planar fusion in simulated cardiac PET.

3 Methods

Dataset We use one noise-free 4D XCAT phantom [9] simulation with ten cardiac phases and paired 3D ground-truth motion vector fields with 1.5-mm isotropic voxel spacing. The image array has shape $(1, 10, 81, 256, 256)$, corresponding to one subject, ten 3D frames, and volumes of $Z \times Y \times X = 81 \times 256 \times 256$. The ground-truth flow array contains nine cumulative reference-to-phase displacement fields from frame 1 to frames $2, \dots, 10$, with shape $(9, 81, 256, 256, 3)$ after reordering to $(Z, Y, X, [d_x, d_y, d_z])$. Pairs $1 \rightarrow 2$ through $1 \rightarrow 9$ are used for model development, while $1 \rightarrow 10$ is reserved as the held-out evaluation pair. Because all pairs share the same anatomy and reference frame and lie on the same simulated periodic cardiac trajectory, this evaluation measures within-phantom

performance rather than generalization across subjects, scanners, or realistic PET count statistics.

Preprocessing PET intensities are scaled to $[0, 1]$, replicated to three channels, and normalized using the ImageNet mean and standard deviation for compatibility with pretrained RAFT. Native 81×256 image planes and corresponding flow fields are bilinearly resized to 128×256 (`align_corners=False`). Because resizing affects only the Z dimension, the dz supervision is scaled by $128/81$, while the horizontal component is unchanged. At inference, predicted dz is multiplied by $81/128$, each component is bilinearly resampled to 81×256 , and the plane is inserted directly into its native volumetric coordinates. No additional axis flip or sign transformation is applied. No padding is required at 128×256 , and boundary slices are processed identically to interior slices.

Multi-view Slice Extraction RAFT predicts dense 2D flow between image pairs. To leverage volumetric ground-truth motion, we construct two orthogonal planar training tasks corresponding to coronal (Y-view) and sagittal (X-view) orientations.

Y-view (coronal). For each fixed index y , we extract 2D planes $I_t^Y(z, x; y)$ and $I_{t'}^Y(z, x; y)$ from volumetric frames at times t and t' , with supervision $\mathbf{f}_Y(z, x; y) = (dx(z, y, x), dz(z, y, x))$.

X-view (sagittal). For each fixed index x , we extract sagittal planes $I_t^X(z, y; x)$ and $I_{t'}^X(z, y; x)$, with supervision $\mathbf{f}_X(z, y; x) = (dy(z, y, x), dz(z, y, x))$.

This formulation provides complementary displacement components while retaining RAFT’s 2D architecture. It assumes that each plane contains sufficient appearance information for its assigned components. Through-plane motion can violate this assumption when structures enter or leave a slice, and the independently trained views have no joint cross-view constraint.

Model and Fine-Tuning Policy RAFT architecture. RAFT [11] takes a pair of $H \times W$ images (here 128×256) and produces a dense 2D displacement field $(u, v) \in \mathbb{R}^2$ at the same resolution. A shared ResNet encoder extracts features at $1/8$ resolution, an all-pairs correlation pyramid is computed between them, and a GRU-based update operator iteratively refines the flow, with the final estimate upsampled via learned convex upsampling. We initialize both view-specific RAFT models from the publicly available `raft-things.pth` checkpoint trained on FlyingThings3D and use 12 update iterations. All model parameters are optimized during fine-tuning, which is performed in FP32. Each RAFT contains 5.26M parameters, giving 10.52M parameters across the two independently fine-tuned models. The bi-planar fusion introduces no additional parameters. On an NVIDIA RTX 3070 Laptop GPU (FP32, batch size 1), the sequential inference pipeline uses approximately 290 MB peak allocated GPU memory and requires 33.93 ± 0.13 s to reconstruct a complete 3D frame pair, including preprocessing, model loading, inverse resampling, volume assembly, and fusion.

Training Protocol Within the development pairs 1→2 through 1→9, individual (phase pair, slice position) samples are randomly split 80/20 into training and validation sets using seed 42. The reserved 1→10 pair is excluded from both subsets. Because the split is not grouped by phase or anatomical position, 14 of 16 slice positions and 7 of 8 phase pairs occur in both subsets for the final $k=16$ configuration. Validation is therefore used for checkpoint selection rather than as an independent estimate of anatomical generalization. Models are optimized using the RAFT sequence loss $\mathcal{L} = \sum_{i=1}^N \gamma^{N-i} \|\mathbf{f}_i - \mathbf{f}_{gt}\|_1$, with $\gamma=0.8$ and $N=12$, using AdamW (lr= 10^{-4} , weight decay 10^{-4} , batch size 2) for 200 epochs.

To examine slice-count sensitivity, we train each view-specific model using $k \in \{4, 8, 16, 32\}$ uniformly spaced slices per volume, corresponding to 1.6%, 3.1%, 6.25%, and 12.5% of the 256 available slice positions. This ablation motivates $k=16$ as a practical accuracy–data tradeoff. To improve performance on the $k=16$ configuration, we then refine slice placement using motion information from the development pairs only. Each anatomical slice is scored by its mean 3D ground-truth displacement magnitude, from which 16 approximately evenly spaced slices are selected from the 64 highest-scoring positions.

Bi-planar Fusion for 3D Motion Reconstruction Given the two trained models, we reconstruct $\hat{\mathbf{u}}(z, y, x) = (\hat{dx}, \hat{dy}, \hat{dz})$ by combining orthogonal predictions on a shared (Z, Y, X) grid. In-plane components are assigned directly: $\hat{dx}(z, y, x) = \hat{f}_{Y,1}(z, x; y)$ from the Y-view and $\hat{dy}(z, y, x) = \hat{f}_{X,1}(z, y; x)$ from the X-view. The shared out-of-plane component is reconciled by averaging:

$$\hat{dz}(z, y, x) = \frac{1}{2}(\hat{f}_{Y,2}(z, x; y) + \hat{f}_{X,2}(z, y; x)).$$

The two dz estimates are averaged only after both views have been mapped back to the common native (Z, Y, X) grid and voxel-displacement units. Equal weighting therefore assumes comparable bias and reliability between views. The independently estimated slice fields are assembled without an explicit inter-slice smoothness, joint cross-view, Jacobian, diffeomorphic, or other deformation-validity constraint. Consequently, low voxelwise error does not by itself guarantee a physically valid deformation field.

Baseline Methods We compare against two reference baselines. **Farneback** uses OpenCV at the native 81×256 slice resolution with per-slice min–max normalization to uint8 and parameters `pyr_scale=0.5`, `levels=5`, `winsize=15`, `iterations=5`, `poly_n=7`, `poly_sigma=1.5`, and `flags=0`. **Pretrained RAFT** uses the same `raft-things.pth` initialization and RAFT preprocessing as the fine-tuned models but without domain adaptation. These baselines quantify improvement relative to classical optical flow and the unadapted RAFT initialization, though they do not constitute a comprehensive comparison with modern 3D deformable-registration methods.

Evaluation Metrics Flow accuracy is evaluated using end-point error (EPE); accuracy below 0.5, 1.0, and 3.0 px; angular error; directional accuracy, defined

as the fraction of locations with angular error below 5° ; and displacement magnitude ratio. For planar magnitude analysis, the ratio of mean predicted to mean ground-truth in-plane magnitude is computed within each motion-containing slice and then averaged across slices. For fused 3D evaluation, magnitude ratio is instead computed once over all pooled moving voxels. The planar and 3D magnitude ratios therefore use different outer aggregations and are not directly comparable. Cross-view dz is evaluated using Pearson correlation together with signed bias and absolute error, since correlation alone does not establish agreement.

4 Results

Training Efficiency We examine how the number of uniformly sampled training slices affects volumetric fusion performance within the single phantom. Figure 1 shows that 3D EPE on the 1→10 evaluation pair improves from 0.535 px at $k=4$ to 0.259 px at $k=16$, with smaller gains beyond 16 slices. We therefore use $k=16$ as a practical accuracy–data tradeoff and apply the motion-rich slice-selection procedure described above for the final model. This ablation characterizes data efficiency within the present phantom rather than a general training requirement.

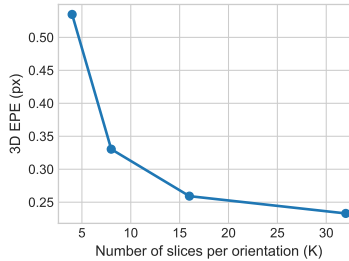


Fig. 1. Within-phantom 3D fusion EPE on the 1→10 evaluation pair vs. number of uniformly sampled training slices per view ($k \in \{4, 8, 16, 32\}$).

Planar Motion Estimation Accuracy Table 1 presents results for both the X-view model, which predicts (dy, dz) motion components, and the Y-view model, which predicts (dx, dz) components. Fine-tuning substantially improves accuracy relative to both Farneback and pretrained RAFT across both views. For the X-view, EPE decreases from 0.696 px to 0.204 px, sub-pixel accuracy increases from 42.3% to 93.0%, and directional accuracy increases from 25.4% to 88.9%. For the Y-view, EPE decreases from 0.514 px to 0.148 px, sub-pixel accuracy reaches 96.9%, and directional accuracy increases from 34.7% to 72.8%.

The fine-tuned Y-view magnitude ratio of 1.600 is an average of per-slice ratios and is sensitive to low-motion slices with small denominators. Over the same set of motion-containing slices, the pooled ratio of mean predicted to mean ground-truth magnitude is 1.165, indicating more moderate aggregate magnitude overestimation despite the larger per-slice-averaged value.

Table 1. Motion estimation accuracy on X-view (sagittal) and Y-view (coronal) slices. Models are evaluated on all slices from the held-out evaluation pair (1→10).

Metric	X-view (Sagittal)			Y-view (Coronal)		
	Farneback	Baseline	Fine-tuned	Farneback	Baseline	Fine-tuned
EPE (px)	0.696	0.449	0.204	0.514	0.318	0.148
Acc <0.5px (%)	42.27	64.91	93.01	65.06	79.87	96.95
Acc <1.0px (%)	77.50	95.16	99.48	87.69	97.06	99.79
Acc <3.0px (%)	99.68	99.81	100.00	98.87	99.92	100.00
Ang. Error (°)	71.78	61.51	16.59	64.38	70.43	29.87
Dir. Accuracy (%)	25.39	38.99	88.90	34.74	26.24	72.80
Mag. Ratio	1.463	0.581	1.172	2.159	0.768	1.600

3D Volumetric Motion Estimation via Bi-Planar Fusion Table 2 compares fine-tuned bi-planar fusion with Farneback and unadapted pretrained-RAFT fusion. Fine-tuned fusion achieves 0.178 px 3D EPE, representing a 74% reduction compared to Farneback (0.676 px) and a 66% reduction compared to pretrained RAFT fusion (0.530 px). Sub-pixel accuracy reaches 96.7%, while directional accuracy improves to 88.6%. The pooled 3D magnitude ratio is 1.033, compared with 1.476 for Farneback fusion and 0.462 for pretrained-RAFT fusion. Thus, over moving voxels in the fused volume, mean predicted displacement magnitude is 3.3% above mean ground-truth magnitude.

Table 2. Volumetric 3D motion estimation accuracy. Metrics are computed voxelwise over the full 3D flow field reconstructed from the held-out evaluation pair (1→10).

Metric	Farneback	Baseline RAFT	Fine-tuned RAFT
EPE (px)	0.676	0.530	0.178
Acc <0.5px (%)	41.74	53.54	96.70
Acc <1.0px (%)	82.42	91.51	99.96
Acc <3.0px (%)	99.48	99.82	100.00
Ang. Error (°)	53.43	61.64	18.13
Dir. Accuracy (%)	51.45	40.11	88.62
Mag. Ratio	1.476	0.462	1.033

Figure 2 shows representative X, Y, and Z planes. Fine-tuned fusion visually follows the principal ground-truth displacement patterns more closely than the

evaluated Farneback and pretrained-RAFT baselines. Because the method imposes no deformation-validity constraint, the qualitative appearance should not be interpreted as evidence of a physically valid or diffeomorphic deformation.

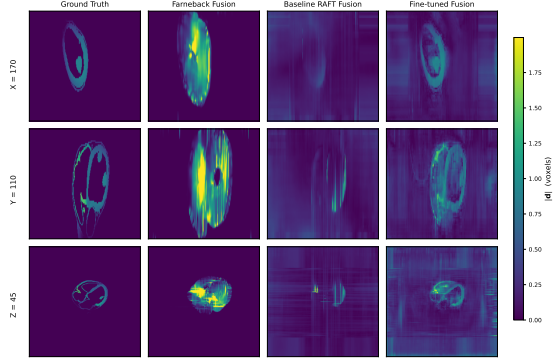


Fig. 2. Qualitative comparison of reconstructed 3D displacement fields on the 1→10 evaluation pair across representative X, Y, and Z planes. Color denotes displacement magnitude on a shared 0–2.0 voxel scale (99.7th percentile; larger values are clipped).

Table 3 summarizes agreement of the independently estimated d_z components with ground truth. The X- and Y-view estimates are strongly correlated with one another ($r=0.880$), but correlation alone does not establish absolute agreement. Their signed biases are nearly identical (-0.0085 and -0.0091 voxels), with a mean cross-view difference of only 0.00055 voxels. Uniform averaging reduces d_z MAE from 0.0648 and 0.0619 voxels to 0.0504 voxels and RMSE from 0.0959 and 0.0909 voxels to 0.0732 voxels, providing empirical support for the fusion rule in this phantom.

Table 3. d_z estimation and fusion on motion voxels of the 1→10 evaluation pair. Bias is mean signed error relative to ground truth.

Metric	Y-view	X-view	Fused
Bias (vox.)	-0.0091	-0.0085	-0.0088
MAE (vox.)	0.0619	0.0648	0.0504
RMSE (vox.)	0.0909	0.0959	0.0732
Corr. with GT	0.925	0.917	0.950

5 Discussion

Bi-planar fusion as a proof of concept Deterministic fusion of two orthogonally fine-tuned 2D models achieves 0.178 px 3D EPE, compared with

0.676 px for Farneback fusion and 0.530 px for unadapted pretrained-RAFT fusion. The similar view-specific d_z biases ($-0.0091/ -0.0085$ voxels) and reduction in d_z MAE to 0.0504 voxels after averaging provide empirical support for equal-weight fusion within this phantom. The pooled 3D magnitude ratio of 1.033 indicates that mean predicted displacement magnitude is 3.3% higher than mean ground-truth magnitude over moving voxels. These results support within-phantom voxelwise motion estimation relative to the evaluated reference baselines, but do not establish deformation validity, cross-subject generalization, or superiority to dedicated 3D registration methods.

Limitations and future validation The present results remain limited to phase pairs from one noise-free XCAT anatomy and therefore do not establish generalization across subjects, abnormal cardiac motion, scanners, or realistic PET count statistics. Validation is additionally correlated across phase and anatomical position, and the final motion-rich training slices are selected using simulated ground-truth motion. Bi-planar reconstruction assumes sufficient within-plane observability under through-plane motion and imposes no explicit inter-slice, joint cross-view, Jacobian, or other deformation-validity constraint. The Farneback and unadapted-RAFT comparisons are reference baselines rather than a comprehensive comparison with modern 3D registration methods, and the measured compute characteristics do not establish an efficiency advantage over such models. Multi-subject simulation, realistic Poisson noise and scanner effects, stronger volumetric baselines, deformation-aware and reliability-aware fusion, and ultimately validation on real gated PET are therefore necessary before claims of clinical readiness.

6 Conclusion

We presented a proof-of-concept bi-planar approach that combines two orthogonally fine-tuned 2D RAFT models into a 3D cardiac displacement field. Within a single noise-free XCAT phantom, fusion achieves 0.178 px 3D EPE (0.27 mm), 96.7% sub-pixel accuracy, and 88.6% directional accuracy, reducing EPE by 74% and 66% relative to the evaluated Farneback and unadapted pretrained-RAFT baselines. The two d_z estimates exhibit comparable bias, and averaging reduces their individual MAE to 0.0504 voxels, providing empirical support for uniform fusion in this phantom. Because training and evaluation share the same simulated anatomy and periodic motion trajectory, these results establish within-phantom feasibility rather than general 3D cardiac-motion or clinical performance. Independent anatomies, realistic PET statistics, stronger volumetric baselines, and real-data validation are necessary next steps.

Disclosure of Interests. The authors have no competing interests to declare that are relevant to the content of this article.

References

1. Balakrishnan, G., Zhao, A., Sabuncu, M.R., Gutttag, J., Dalca, A.V.: Voxelmorph: A learning framework for deformable medical image registration. *IEEE Transactions on Medical Imaging* **38**(8), 1788–1800 (2019). <https://doi.org/10.1109/TMI.2019.2897538>
2. Dawood, M., Gigengack, F., Jiang, X., Schäfers, K.P.: A mass conservation-based optical flow method for cardiac motion correction in 3D-PET. *Medical Physics* **40**(1), 012505 (2013). <https://doi.org/10.1118/1.4770276>
3. Dong, G., Dai, J., Li, N., Zhang, C., He, W., Liu, L., Chan, Y., Li, Y., Xie, Y., Liang, X.: 2d/3d non-rigid image registration via two orthogonal x-ray projection images for lung tumor tracking. *Bioengineering* **10**(2), 144 (2023). <https://doi.org/10.3390/bioengineering10020144>
4. Dosovitskiy, A., Fischer, P., Ilg, E., Häusser, P., Hazirbas, C., Golkov, V., Smagt, P.v.d., Cremers, D., Brox, T.: FlowNet: Learning optical flow with convolutional networks. In: 2015 IEEE International Conference on Computer Vision (ICCV). pp. 2758–2766 (2015). <https://doi.org/10.1109/ICCV.2015.316>
5. Hong, I., Jones, J., Casey, M.: Ultrafast elastic motion correction via motion deblurring. In: 2014 IEEE Nuclear Science Symposium and Medical Imaging Conference (NSS/MIC). pp. 1–2 (2014). <https://doi.org/10.1109/NSSMIC.2014.7430841>
6. Horn, B.K.P., Schunck, B.G.: Determining optical flow. *Artificial Intelligence* **17**(1), 185–203 (1981). [https://doi.org/10.1016/0004-3702\(81\)90024-2](https://doi.org/10.1016/0004-3702(81)90024-2)
7. Masutani, E.M., Chandrupatla, R.S., Wang, S., Zocchi, C., Hahn, L.D., Horowitz, M., Jacobs, K., Kligerman, S., Raimondi, F., Patel, A., Hsiao, A.: Deep learning synthetic strain: Quantitative assessment of regional myocardial wall motion at MRI. *Radiology: Cardiothoracic Imaging* **5**(3), e220202 (2023). <https://doi.org/10.1148/ryct.220202>
8. Petibon, Y., Sun, T., Han, P.K., Ma, C., El Fakhri, G., Ouyang, J.: MR-based cardiac and respiratory motion correction of PET: application to static and dynamic cardiac ^{18}F -FDG imaging. *Physics in Medicine & Biology* **64**(19), 195009 (2019). <https://doi.org/10.1088/1361-6560/ab39c2>
9. Segars, W.P., Sturgeon, G., Mendonca, S., Grimes, J., Tsui, B.M.W.: 4D XCAT phantom for multimodality imaging research. *Medical Physics* **37**(9), 4902–4915 (2010). <https://doi.org/10.1118/1.3480985>
10. Sun, D., Yang, X., Liu, M.Y., Kautz, J.: PWC-Net: CNNs for optical flow using pyramid, warping, and cost volume. In: 2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 8934–8943. IEEE Computer Society, Los Alamitos, CA, USA (Jun 2018). <https://doi.org/10.1109/CVPR.2018.00931>
11. Teed, Z., Deng, J.: RAFT: Recurrent all-pairs field transforms for optical flow. In: Vedaldi, A., Bischof, H., Brox, T., Frahm, J.M. (eds.) *Computer Vision – ECCV 2020*. pp. 402–419. Springer International Publishing, Cham (2020). https://doi.org/10.1007/978-3-030-58536-5_24
12. Van Houtte, J., Audenaert, E., Zheng, G., Sijbers, J.: Deep learning-based 2d/3d registration of an atlas to biplanar x-ray images. *International Journal of Computer Assisted Radiology and Surgery* **17**(7), 1333–1342 (2022). <https://doi.org/10.1007/s11548-022-02586-3>
13. Zhang, Y., Liao, Q., Ding, L., Zhang, J.: Bridging 2d and 3d segmentation networks for computation-efficient volumetric medical image segmentation: An empirical study of 2.5d solutions. *Computerized Medical Imaging and Graphics* **99**, 102088 (2022). <https://doi.org/10.1016/j.compmedimag.2022.102088>

14. Zou, S.J., Chinn, G., Levin, C.S.: Motion correction for cardiac PET via parametric image optical flow. In: 2024 IEEE Nuclear Science Symposium (NSS), Medical Imaging Conference (MIC) and Room Temperature Semiconductor Detector Conference (RTSD). pp. 1–1 (2024). <https://doi.org/10.1109/NSS/MIC/RTSD57108.2024.10657995>