

Vis2Reg: Visibility-Aware Landmark-Free Geometric 3D–2D Registration for Liver Laparoscopy

Jiaming Feng¹[0009–0009–0506–7152], Xukun Zhang²[0000–0003–2869–9434], Shahid Farid³[0000–0001–7685–182X], and Sharib Ali(✉)¹[0000–0003–1313–3542]

- ¹ AI in Medicine and Surgery Group, School of Computer Science, University of Leeds, LS2 9JT, Leeds, United Kingdom
{vqkc0507, s.s.ali}@leeds.ac.uk
- ² Department of Diagnostic Radiology, Li Ka Shing Faculty of Medicine, The University of Hong Kong, Pok Fu Lam, Hong Kong
xukunzhg@hku.hk
- ³ Department of HPB and Transplant Surgery, St. James’s University Hospital, Leeds, United Kingdom
s.farid@nhs.net

Abstract. Accurate 3D–2D liver registration, which aligns preoperative 3D models to partial, view-dependent intraoperative surface observations, is critical for AR-guided laparoscopic surgery but remains challenging due to severe occlusion, limited visibility, and the lack of 3D ground-truth supervision. Existing landmark-free approaches perform partial-to-complete geometric alignment, yet robust self-supervision under extreme partial visibility remains difficult. We propose Vis2Reg, a visibility-aware registration framework that explicitly constrains deformation using mask-consistent visible regions. We introduce a visibility-aware self-supervision that derives a visible-domain 3D supervision signal from intraoperative masks, enabled by differentiable point rasterization and mask-guided back-projection. This formulation improves robustness under severe occlusion while maintaining fully self-supervised learning. Vis2Reg combines a robust geometric rigid initialization module with an implicit neural deformation field for stable alignment. Vis2Reg achieves a Dice score of 92.6% and a Chamfer Distance of 1.43 mm on real intraoperative datasets, with 111 ms per-frame inference time, demonstrating both accuracy and practical efficiency. The source code is publicly available at <https://github.com/aimsgroup-Leeds/Vis2Reg.git>.

Keywords: 3D–2D Registration · Laparoscopic Liver Surgery · Visibility-Aware Self-Supervised Learning

1 Introduction

Minimally invasive laparoscopic and robotic liver resection is increasingly adopted, but surgeons rely on a narrow field-of-view with no direct access to sub-surface

anatomy, limiting assessment of tumour–vessel relationships and increasing surgical risk [5,1]. Augmented reality (AR) can mitigate this by overlaying patient-specific preoperative 3D liver, vascular, or tumour models onto the intraoperative video, improving spatial awareness beyond the visible surface [1,17,21]. Accurate, robust preoperative-to-intraoperative fusion of the deformable liver is therefore critical for reliable AR guidance.

At its core, AR-guided laparoscopy requires preoperative-to-intraoperative 3D–2D registration of a preoperative liver model to intraoperative views and camera geometry, which is ill-posed under large non-rigid deformation, adverse imaging (occlusion, specularities, smoke, blood), and partial, view-dependent observations [4,12,22,10,11]. Early approaches relied on landmarks, contours, or biomechanical constraints [4,2,13], but are brittle under occlusion and low overlap [11,22,13]. This family spans classical iterative closest point-based (ICP) optimization, biomechanical finite element method (FEM) simulation, and shape-matching such as functional maps [2,16,15]. Learning-based approaches predict deformation directly from observations [12] yet remain limited by scarce supervision and incomplete views [1,13].

A recent paradigm shift reformulates laparoscopic registration as *partial-to-complete 3D registration*, aligning sparse intraoperative surface observations with the full preoperative liver geometry [26,7,14]. This formulation enables self-supervised learning using readily available signals such as segmentation masks and geometric consistency, eliminating the need for explicit 3D correspondence supervision [26,7]. Nevertheless, intraoperative supervision is fundamentally visibility limited: only the mask-consistent visible surface provides reliable geometric constraints, and explicitly modeling such visibility-consistent supervision to guide deformation remains challenging under severe occlusion and limited overlap [6,3,22,11,14]. Furthermore, rigid initialization is often unstable when intraoperative observations are sparse and noisy [22,10,7,14]. Generic low-overlap point-cloud registration methods (e.g., PREDATOR [8], GeoTransformer [19]) improve correspondence robustness but are not tailored to visibility-limited surgical self-supervision. To address these challenges, we propose *Vis2Reg*, a visibility-aware, landmark-free, geometry-only framework for self-supervised laparoscopic liver registration under severe partial visibility. It derives mask-consistent visible-domain 3D supervision from intraoperative masks via differentiable point rasterization and mask-guided back-projection [9], focusing learning on truly observable geometry. This improves deformation robustness while preserving anatomical plausibility, and enables geometry-only inference without masks or explicit correspondences.

Our contributions are three-fold: (1) We introduce a visibility-aware self-supervision formulation that constructs mask-consistent visible-domain 3D supervision signals for landmark-free laparoscopic liver registration; (2) we develop a differentiable point rasterization and mask-guided back-projection mechanism to derive visibility-consistent geometric supervision from intraoperative observations; and (3) we integrate robust geometric rigid initialization with an implicit

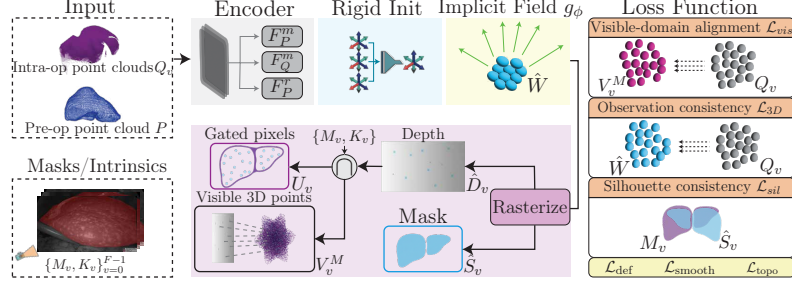


Fig. 1. Vis2Reg pipeline. Given preoperative P and intraoperative partial Q , the network outputs warped geometry $\hat{W}$ via rigid initialization and non-rigid deformation; differentiable rasterization and mask-gated back-projection construct the visible-domain supervision set V_v^M from $(\hat{D}_v, \hat{S}_v)$ and M_v .

neural deformation field, achieving improved registration robustness and near-real-time performance on in-vivo laparoscopic datasets.

2 Method

2.1 Problem formulation

Vis2Reg addresses self-supervised intraoperative liver registration under severe partial visibility and without 3D-2D ground-truth supervision. The clinical goal is a 3D-to-2D AR overlay, achieved as partial-to-complete 3D registration aligning the complete preoperative cloud P with the partial intraoperative observation Q . Each view Q_v is reconstructed from laparoscopic images via monocular depth (DepthAnything [24]), the liver mask M_v , and back-projection with intrinsics K_v . Given a preoperative liver point cloud $P = \{p_i\}_{i=1}^N$ and F view-dependent intraoperative point clouds $\{Q_v\}_{v=0}^{F-1}$ (all in camera coordinates), together with F -view masks $\{M_v\}_{v=0}^{F-1}$ and intrinsics $\{K_v\}_{v=0}^{F-1}$, we denote the merged intraoperative observation as $Q \triangleq \bigcup_{v=0}^{F-1} Q_v$. We use $F=3$ as a short *local, non-temporal* window corresponding to a near-static liver; M_v and K_v serve as per-view supervision signals, not model inputs. The objective is to estimate a rigid alignment and a non-rigid deformation modeled as an implicit displacement field g_ϕ , producing a warped point cloud $\hat{W}$ that aligns P to the intraoperative observations under visibility-consistent geometric constraints. As shown in Fig. 1, this is achieved using a geometry-only registration network comprising a robust rigid initialization module and the implicit deformation field g_ϕ . During training, visibility-consistent supervision is constructed from masks via differentiable rasterization technique and mask-gated back-projection, while at inference the model operates using geometric inputs only.

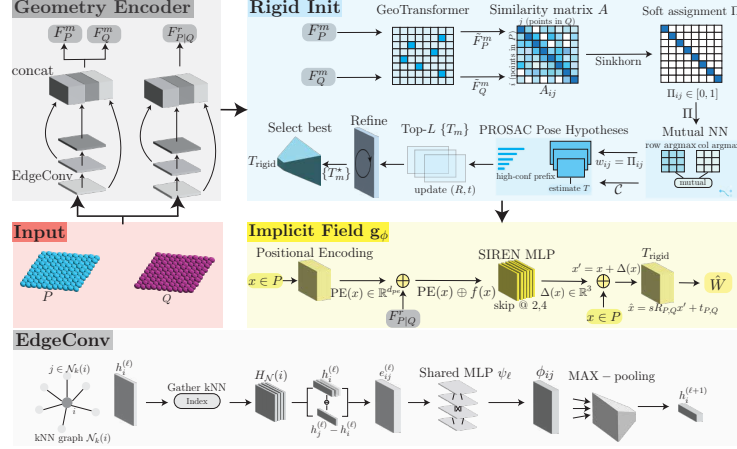


Fig. 2. Architecture of the geometry-only registration network in Vis2Reg: EdgeConv feature encoding, rigid initialization (T_{rigid}), and implicit non-rigid deformation field g_ϕ producing the warped point cloud $\hat{W}$. Bottom: EdgeConv update block.

2.2 Geometry-Only Registration Network

As illustrated in Fig. 1, Vis2Reg predicts the warped geometry $\hat{W}$ via (i) geometric feature encoders, (ii) a robust rigid initialization module for stable global alignment under partial visibility, and (iii) an implicit displacement field g_ϕ modeling continuous local deformation. Concretely, the rigid module builds on a GeoTransformer matcher and the deformation field is implemented as a multilayer perceptron (MLP) with sinusoidal representation network (SIREN) activations, both detailed below.

Geometric Feature Encoders. In Fig. 2, we extract multi-scale geometric features using a multi-layer EdgeConv encoder [23], which captures local shape context through dynamic neighborhood aggregation. Given a point set $X = \{x_i\}_{i=1}^n$ with initial features $h_i^{(0)} = x_i$, each layer updates point features via

$$h_i^{(\ell+1)} = \max_{j \in \mathcal{N}_k(i)} \psi_\ell \left(\left[h_i^{(\ell)} \parallel h_j^{(\ell)} - h_i^{(\ell)} \right] \right), \quad (1)$$

where $\psi_\ell(\cdot)$ is a shared MLP and $\mathcal{N}_k(i)$ denotes the k -nearest-neighbor (k -NN) neighborhood. Concatenating intermediate outputs yields multi-scale features $F(X)$. We employ two encoders with distinct roles: a Siamese *matching encoder* $\mathcal{E}_m$ that extracts correspondence features for rigid alignment from (P, Q) , while a *registration encoder* $\mathcal{E}_r$ extracts deformation features conditioned on the pair (P, Q) for the implicit field, i.e., $F_P^m = \mathcal{E}_m(P)$, $F_Q^m = \mathcal{E}_m(Q)$, and $F_{P|Q}^r = \mathcal{E}_r(P, Q)$, queried at the source points P .

Rigid Seed Initialization. To obtain a stable rigid alignment under sparse, noisy observations, we estimate an initial pose via soft correspondence, hypothesis evaluation, and geometric refinement (Fig. 2). GeoTransformer contextualizes

F_P^m, F_Q^m into $\tilde{F}_P^m, \tilde{F}_Q^m$, giving a similarity matrix $A_{ij} = \langle \tilde{f}_i, \tilde{g}_j \rangle / \tau$; Sinkhorn normalization yields a soft assignment Π with confidence weights $w_{ij} = \Pi_{ij}$, and mutual nearest neighbors form a correspondence set $\mathcal{C}$. Progressive Sample Consensus (PROSAC) then generates pose hypotheses $\{T_m\}$ from minimal samples in $\mathcal{C}$, ranked by weighted inlier scores; the top hypotheses are refined by trimmed point-to-point and point-to-plane ICP, and the best provides the rigid initialization for g_ϕ .

Implicit Non-Rigid Deformation Field. We represent non-rigid deformation as a continuous implicit displacement field g_ϕ conditioned on geometric features. For each source point $x \in P$, positional encoding $\text{PE}(x) \in \mathbb{R}^{d_{pe}}$ is concatenated with its pair-conditioned registration feature $f(x) = F_{P|Q}^r(x)$ to form the field input $z(x) = \text{PE}(x) \oplus f(x)$. A SIREN-based MLP models the implicit field and predicts a displacement $\Delta(x) = g_\phi(z(x)) \in \mathbb{R}^3$, which produces a locally deformed point $x' = x + \Delta(x)$. The final warped geometry is obtained by applying the rigid alignment to the deformed point:

$$\hat{x} = sR_{P,Q} x' + t_{P,Q}, \quad \hat{W} = \{\hat{x} \mid x \in P\}. \quad (2)$$

The resulting warped point cloud $\hat{W}$ is then used for visibility-aware supervision. Since both the pair-conditioned feature $F_{P|Q}^r$ and the rigid transform $(sR_{P,Q}, t_{P,Q})$ depend on Q , the deformation adapts to each intraoperative input rather than being a fixed function of P .

2.3 Visibility-Aware Self-Supervised Learning

The observed clouds $\{Q_v\}_{v=0}^{F-1}$ are view-dependent visible subsets of the surface, with no 3D registration ground truth. Vis2Reg therefore constructs a *mask-consistent visible-domain 3D supervision signal*, restricting supervision to geometrically observable regions defined by the masks and guiding deformation only by physically valid observations. Unlike the symmetric rendered-mask consistency of Self-P2IR [26], our supervision is mask-gated and one-way (observation-to-model), so unobserved model regions are never penalized.

Visible-domain construction. Let Ω denote the discrete $W \times H$ image lattice. From the warped cloud $\hat{W}$, intrinsics $\{K_v\}$, and masks $\{M_v\}$, we render depth and silhouette per view by differentiable point rasterization:

$$(\hat{S}_v, \hat{D}_v) = \mathcal{R}_{K_v}(\hat{W}), \quad v = 0, \dots, F-1, \quad (3)$$

where $\mathcal{R}_{K_v}$ is parameterized by K_v and $\hat{D}_v(u) = 0$ marks empty pixels. We define the mask-consistent visible pixel set $U_v = \{u \in \Omega \mid \hat{D}_v(u) > 0 \wedge M_v(u) = 1\}$ and back-project it to an explicit visible-domain 3D supervision set:

$$V_v^M = \left\{ \pi_{K_v}^{-1}(u, \hat{D}_v(u)) \mid u \in U_v \right\}, \quad (4)$$

where $\pi_{K_v}^{-1}$ is the analytic back-projection from a pixel-depth pair to a camera-frame 3D point.

Self-Supervised Objectives. We optimize a weighted combination of visibility-aware alignment losses and deformation regularization:

$$\mathcal{L} = \lambda_{3D}\mathcal{L}_{3D} + \lambda_{vis}\mathcal{L}_{vis} + \lambda_{sil}\mathcal{L}_{sil} + \lambda_{def}\mathcal{L}_{def} + \lambda_{smooth}\mathcal{L}_{smooth} + \lambda_{topo}\mathcal{L}_{topo}. \quad (5)$$

We denote the one-way Chamfer as $\text{CD}_{\rightarrow}(A, B) = \frac{1}{|A|} \sum_{a \in A} \min_{b \in B} \|a - b\|_2^2$ and the symmetric Chamfer as $\text{CD}(A, B) = \text{CD}_{\rightarrow}(A, B) + \text{CD}_{\rightarrow}(B, A)$.

The visibility-aware alignment terms are an observation-to-model one-way Chamfer $\mathcal{L}_{3D} = \frac{1}{F} \sum_v \text{CD}_{\rightarrow}(Q_v, \hat{W})$ (handling the partial visibility of Q_v), a symmetric visible-domain term $\mathcal{L}_{vis} = \frac{1}{F} \sum_v \text{CD}(V_v^M, Q_v)$ aligning the masked set V_v^M with Q_v , and a silhouette term $\mathcal{L}_{sil} = \frac{1}{F} \sum_v [\text{BCE}(\hat{S}_v, M_v) + \text{DiceLoss}(\hat{S}_v, M_v)]$. The displacement field is regularized by deformation magnitude ($\mathcal{L}_{def} = \frac{1}{N} \sum_i \|\Delta_i\|_2^2$), local smoothness ($\mathcal{L}_{smooth} = \frac{1}{N} \sum_i \frac{1}{|\mathcal{N}(i)|} \sum_{j \in \mathcal{N}(i)} \|\Delta_i - \Delta_j\|_2^2$), and local distance or topology preservation ($\mathcal{L}_{topo} = \frac{1}{N} \sum_i \frac{1}{|\mathcal{N}(i)|} \sum_{j \in \mathcal{N}(i)} (\|\hat{x}_i - \hat{x}_j\|_2 - \|x_i - x_j\|_2)^2$). Here, $\mathcal{N}(i)$ is the k -NN neighborhood of x_i and $\Delta_i = \Delta(x_i)$.

Training schedule. We adopt a two-stage synthetic→real schedule: the matching and rigid-seed modules are pretrained on synthetic data with ground-truth (GT) poses (non-rigid field disabled), then trained on real data (no GT) with a short rigid warm-up followed by joint visibility-aware optimization.

3 Experiments

3.1 Dataset and Implementation Details

We evaluate Vis2Reg on the public P2I-LReg dataset, a real intraoperative liver registration benchmark with 346 keyframes from 21 patients [26]. Each sample provides a preoperative 3D liver model, an intraoperative sparse point cloud, a binary segmentation mask, and camera intrinsics. Following the official patient-level protocol, patients are split into 12/4/5 for training/validation/testing, and results are averaged over 5-fold cross-validation. For robust rigid initialization, we additionally use the patient-specific Landmark-Free Synthetic Dataset (red-green-blue plus depth (RGB-D) observations rendered from preoperative models via physics-based Blender simulation with diverse viewpoints and occlusions; 2500 samples per patient; 60%, 20%, and 20% split; training-fold patients only).

Inputs are P , Q , and multi-view masks and intrinsics $\{M_v, K_v\}$, all in camera coordinates at real-world scale (no centering and normalization). During training we apply random dropout and jittering to P and statistical denoising to Q ; both are resampled or zero-padded to $N_{\max} = 6000$.

We use AdamW (lr = 3×10^{-4} , wd = 10^{-4}) with cosine scheduling and automatic mixed precision (AMP). In Stage-2 real-data training, we train for 60 epochs (batch size 1) with a 20-epoch rigid warm-up (non-rigid branch and deformation regularizers frozen), followed by joint optimization; a 3-frame multi-view input is used by default. Loss weights are $\lambda_{3D} = 0.5$, $\lambda_{vis} = 1.2$, $\lambda_{sil} = 1.0$ (BCE

Table 1. Comparison on synthetic rigid initialization and real intraoperative registration (mean $\pm$ std). RRE and RTE are reported only for rigid-comparable methods; non-rigid methods are marked as “—” in the synthetic block.

Method	Synthetic (Rigid Init)		Real Intraoperative (Final)	
	RRE ($^{\circ}$) $\downarrow$	RTE (mm) $\downarrow$	Dice (%) $\uparrow$	CD (mm) $\downarrow$
GeoTransformer [19]	31.30	78.10	71.16 ± 6.46	$3.43 \pm \mathbf{0.96}$
DPF [18]	—	—	$72.19 \pm \mathbf{5.97}$	3.31 ± 1.18
PointSetReg [25]	—	—	$75.24 \pm \mathbf{5.92}$	$3.20 \pm \mathbf{1.03}$
Self-P2IR [26]	0.21	1.32	78.89 ± 6.76	2.97 ± 1.06
Ours (Vis2Reg)	0.08	0.26	92.60 ± 7.24	1.43 ± 1.26

and Dice weights both 1.0), $\lambda_{def} = 0.1$, $\lambda_{smooth} = 0.1$, and $\lambda_{topo} = 0.3$. For differentiable rasterization, *points-per-pixel*=16 and the maximum number of visible points is 5000. On synthetic data, we report relative rotation error (RRE) and relative translation error (RTE), defined as $RRE = \arccos((\text{tr}(\mathbf{R}_{\text{pred}}^T \mathbf{R}_{\text{gt}}) - 1)/2)$ (degrees) and $RTE = \|\mathbf{t}_{\text{pred}} - \mathbf{t}_{\text{gt}}\|_2$ (mm). On real data, we report Chamfer Distance ($\mathcal{L}_{3D}$) and silhouette Dice. Dice (our primary AR-overlay metric) measures overlap between the rasterized registered silhouette and the liver mask, whereas CD is a complementary surface-distance measure, not a target registration error. All baselines use the same patient-level folds and protocol implemented in PyTorch on NVIDIA L40S GPU.

3.2 Comparison and Ablation Study

On the synthetic benchmark (with 3D GT), Vis2Reg achieves the best rigid initialization performance in the synthetic block of Table 1, with RRE 0.08° and RTE 0.26 mm, outperforming GeoTransformer (RRE 31.30° , RTE 78.10 mm) and Self-P2IR (RRE 0.21° , RTE 1.32 mm), indicating a more reliable rigid seed under low overlap. DPF and PointSetReg are non-rigid models and are therefore not included in the synthetic rigid comparison. The GeoTransformer entry uses the standalone matcher with its built-in closed-form pose estimation, with no robust outlier rejection (e.g., RANSAC) and no ICP refinement, hence the higher RRE/RTE (Table 1). Vis2Reg instead uses the full pipeline (GeoTransformer matching, mutual filtering, PROSAC, trimmed ICP), reconciling the gap with the robust-estimator values reported on the source dataset [26]. On the real intraoperative benchmark (Table 1), Vis2Reg achieves the best mean performance in both final registration metrics, with Dice $92.60 \pm 7.24\%$ and CD 1.43 ± 1.26 mm. Compared with the strongest prior method (Self-P2IR), Vis2Reg improves Dice by 13.71 pp (percentage points) and reduces CD by 1.54 mm. These gains are consistent with our visibility-aware design, where mask-gated visible-domain supervision constrains optimization to observed regions and mitigates drift under severe occlusion and partial visibility. Importantly, Self-P2IR already attains a near-perfect rigid init (RRE 0.21° , RTE 1.32 mm), comparable to ours, yet trails Vis2Reg by 13.71 pp in Dice; as both start from a strong rigid seed, this gain stems from the visibility-aware non-rigid supervision, not initialization. The

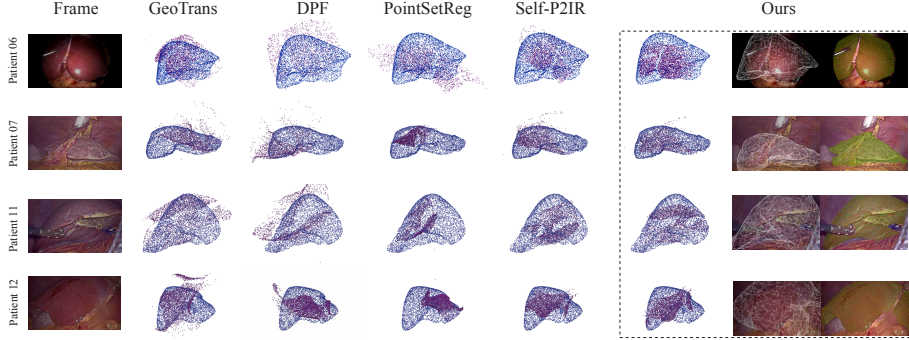


Fig. 3. Qualitative comparison on intraoperative keyframes. Col. 1: input. Cols. 2–5: baselines. Col. 6: Vis2Reg point cloud. Col. 7: Vis2Reg mesh overlay. Col. 8: visible-region fitting. Blue: preoperative; purple: intraoperative.

Table 2. Ablations on real data. “With gate” refers to $U_v = \{\hat{D}_v > 0 \wedge M_v = 1\}$, while “without gate” uses $U_v = \{\hat{D}_v > 0\}$.

Variant	$\mathcal{L}_{vis}$	gate	1-way $\mathcal{L}_{3D}$	seed	Dice $\uparrow$	CD $\downarrow$ (mm)
w/o $\mathcal{L}_{vis}$	N	Y	Y	Y	74.29	3.12
w/o mask-gating	Y	N	Y	Y	79.57	3.04
sym. $\mathcal{L}_{3D}$	Y	Y	N	Y	<u>84.48</u>	<u>2.98</u>
weak rigid init	Y	Y	Y	N	69.32	3.64
Full (Vis2Reg)	Y	Y	Y	Y	92.60	1.43

ablations (Table 2) corroborate this: removing visibility-aware supervision ($\mathcal{L}_{vis}$ or mask-gating) alone reduces Dice by 13–18 pp.

Fig. 3 shows qualitative results on representative keyframes: Vis2Reg achieves tighter preoperative–intraoperative alignment than the baselines and better adapts the visible surface to observations while preserving globally plausible anatomy. Vis2Reg runs at 111.38 ms per forward pass with 1.95 GB GPU memory on an NVIDIA L40S, supporting near-real-time intraoperative use. The reported inference time covers forward registration after Q_v and M_v are available, excluding depth and mask reconstruction.

Ablation results in Table 2 validate each design choice: removing $\mathcal{L}_{vis}$ or mask-gating, replacing one-way $\mathcal{L}_{3D}$ with symmetric Chamfer, and weakening rigid initialization each degrade performance—the last causing the largest drop—supporting mask-gated visible-domain supervision, the observation-to-model formulation, and a strong rigid seed.

4 Conclusion

Vis2Reg is a visibility-aware, landmark-free, geometry-only framework for intraoperative liver 3D–2D registration under severe occlusion, coupling robust rigid

initialization with an implicit non-rigid deformation field and differentiable point rasterization for mask- and camera-based self-supervision without 3D ground truth. On real intraoperative data, it significantly improves Dice and Chamfer Distance over prior methods while retaining near-real-time performance. We restrict our evaluation and claims to the P2I-LReg benchmark. Evaluating tumour and vascular structures would require internal anatomical ground truth [20], unavailable here, and is left to future work.

Acknowledgments. The project was funded by the Engineering and Physical Sciences Research Council [Grant No. UKRI914].

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

References

1. Ali, S., Espinel, Y., Jin, Y., Liu, P., Güttner, B., Zhang, X., Zhang, L., Dowrick, T., Clarkson, M.J., Xiao, S., Wu, Y., Yang, Y., Zhu, L., Sun, D., Li, L., Pfeiffer, M., Farid, S., Maier-Hein, L., Buc, E., Bartoli, A.: An objective comparison of methods for augmented reality in laparoscopic liver resection by preoperative-to-intraoperative image fusion from the MICCAI 2022 challenge. *Medical Image Analysis* **99**, 103371 (2025). <https://doi.org/10.1016/j.media.2024.103371>
2. Besl, P.J., McKay, N.D.: A method for registration of 3D shapes. *IEEE Transactions on Pattern Analysis and Machine Intelligence* **14**(2), 239–256 (1992). <https://doi.org/10.1109/34.121791>
3. Dai, Y., Yang, X., Hao, J., Luo, H., Mei, G., Jia, F.: Preoperative and intraoperative laparoscopic liver surface registration using deep graph matching of representative overlapping points. *International Journal of Computer Assisted Radiology and Surgery* **20**(2), 269–278 (2025). <https://doi.org/10.1007/s11548-024-03312-x>
4. Espinel, Y., Calvet, L., Botros, K., Buc, E., Tilmant, C., Bartoli, A.: Using multiple images and contours for deformable 3D-2D registration of a preoperative CT in laparoscopic liver surgery. *International Journal of Computer Assisted Radiology and Surgery* **17**(12), 2211–2219 (2022). <https://doi.org/10.1007/s11548-022-02774-1>
5. Feuerstein, M., Mussack, T., Heining, S.M., Navab, N.: Intraoperative laparoscope augmentation for port placement and resection planning in minimally invasive liver resection. *IEEE Transactions on Medical Imaging* **27**(3), 355–369 (2008). <https://doi.org/10.1109/TMI.2007.907327>
6. Guan, P., Luo, H., Guo, J., Zhang, Y., Jia, F.: Intraoperative laparoscopic liver surface registration with preoperative CT using mixing features and overlapping region masks. *International Journal of Computer Assisted Radiology and Surgery* **18**(8), 1521–1531 (2023). <https://doi.org/10.1007/s11548-023-02846-w>
7. Huang, B., Yang, X., Jia, F.: A landmark-free 3D–2D rigid liver registration via point cloud matching for laparoscopic surgery. *Healthcare Technology Letters* **12**(1), e70030 (2025). <https://doi.org/10.1049/htl2.70030>
8. Huang, S., Gojcic, Z., Usvyatsov, M., Wieser, A., Schindler, K.: PREDATOR: Registration of 3D point clouds with low overlap. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. pp. 4267–4276 (2021). <https://doi.org/10.1109/CVPR46437.2021.00425>

9. Johnson, J., Ravi, N., Reizenstein, J., Novotny, D., Tulsiani, S., Lassner, C., Branson, S.: Accelerating 3D deep learning with PyTorch3D. In: SIGGRAPH Asia 2020 Courses. p. 10. Association for Computing Machinery (2020). <https://doi.org/10.1145/3415263.3419160>
10. Koo, B., Robu, M.R., Allam, M., Pfeiffer, M., Thompson, S., Gurusamy, K., Davidson, B., Speidel, S., Hawkes, D., Stoyanov, D., Clarkson, M.J.: Automatic, global registration in laparoscopic liver surgery. *International Journal of Computer Assisted Radiology and Surgery* **17**(1), 167–176 (2022). <https://doi.org/10.1007/s11548-021-02518-7>
11. Maier-Hein, L., Mountney, P., Bartoli, A., Elhawary, H., Elson, D., Groch, A., Kolb, A., Rodrigues, M., Sorger, J., Speidel, S., Stoyanov, D.: Optical techniques for 3D surface reconstruction in computer-assisted laparoscopic surgery. *Medical Image Analysis* **17**(8), 974–996 (2013). <https://doi.org/10.1016/j.media.2013.04.003>
12. Mhiri, I., Pizarro, D., Bartoli, A.: Neural patient-specific 3D-2D registration in laparoscopic liver resection. *International Journal of Computer Assisted Radiology and Surgery* **20**(1), 57–64 (2025). <https://doi.org/10.1007/s11548-024-03231-x>
13. Neri, A., Penza, V., Baldini, C., Mattos, L.S.: Surgical augmented reality registration methods: A review from traditional to deep learning approaches. *Computerized Medical Imaging and Graphics* **124**, 102616 (2025). <https://doi.org/10.1016/j.compmedimag.2025.102616>
14. Neri, A., Penza, V., Haouchine, N., Mattos, L.S.: Benchmarking complete-to-partial point cloud registration techniques for laparoscopic surgery. *Frontiers in Robotics and AI* **12**, 1702360 (2025). <https://doi.org/10.3389/frobt.2025.1702360>
15. Ovsjanikov, M., Ben-Chen, M., Solomon, J., Butscher, A., Guibas, L.: Functional maps: A flexible representation of maps between shapes. *ACM Transactions on Graphics* **31**(4), 30:1–30:11 (2012). <https://doi.org/10.1145/2185520.2185526>
16. Özgür, E., Koo, B., Le Roy, B., Buc, E., Bartoli, A.: Preoperative liver registration for augmented monocular laparoscopy using backward-forward biomechanical simulation. *International Journal of Computer Assisted Radiology and Surgery* **13**(10), 1629–1640 (2018). <https://doi.org/10.1007/s11548-018-1842-3>
17. Prevost, G.A., Eigl, B., Paolucci, I., Rudolph, T., Peterhans, M., Weber, S., Beldi, G., Candinas, D., Lachenmayer, A.: Efficiency, accuracy and clinical applicability of a new image-guided surgery system in 3D laparoscopic liver surgery. *Journal of Gastrointestinal Surgery* **24**(10), 2251–2258 (2020). <https://doi.org/10.1007/s11605-019-04395-7>
18. Prokudin, S., Ma, Q., Raafat, M., Valentin, J., Tang, S.: Dynamic point fields. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*. pp. 7964–7976 (2023)
19. Qin, Z., Yu, H., Wang, C., Guo, Y., Peng, Y., Ilic, S., Hu, D., Xu, K.: Geo-Transformer: Fast and robust point cloud registration with geometric transformer. *IEEE Transactions on Pattern Analysis and Machine Intelligence* **45**(8), 9806–9821 (2023). <https://doi.org/10.1109/TPAMI.2023.3259038>
20. Rabbani, N., Calvet, L., Espinel, Y., Le Roy, B., Ribeiro, M., Buc, E., Bartoli, A.: A methodology and clinical dataset with ground-truth to evaluate registration accuracy quantitatively in computer-assisted laparoscopic liver resection. *Computer Methods in Biomechanics and Biomedical Engineering: Imaging & Visualization* **10**(4), 441–450 (2022). <https://doi.org/10.1080/21681163.2021.2002195>
21. Ramalhinho, J., Yoo, S., Dowrick, T., Koo, B., Somasundaram, M., Gurusamy, K., Hawkes, D.J., Davidson, B., Blandford, A., Clarkson, M.J.: The value of Augmented Reality in surgery—a usability study on laparoscopic liver surgery. *Medical Image Analysis* **90**, 102943 (2023). <https://doi.org/10.1016/j.media.2023.102943>

22. Robu, M.R., Ramalhinho, J., Thompson, S., Gurusamy, K., Davidson, B., Hawkes, D., Stoyanov, D., Clarkson, M.J.: Global rigid registration of CT to video in laparoscopic liver surgery. *International Journal of Computer Assisted Radiology and Surgery* **13**(6), 947–956 (2018). <https://doi.org/10.1007/s11548-018-1781-z>
23. Wang, Y., Sun, Y., Liu, Z., Sarma, S.E., Bronstein, M.M., Solomon, J.M.: Dynamic graph CNN for learning on point clouds. *ACM Transactions on Graphics* **38**(5), 146:1–146:12 (2019). <https://doi.org/10.1145/3326362>
24. Yang, L., Kang, B., Huang, Z., Xu, X., Feng, J., Zhao, H.: Depth anything: Unleashing the power of large-scale unlabeled data. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. pp. 10371–10381 (2024). <https://doi.org/10.1109/CVPR52733.2024.00987>
25. Zhao, M., Jiang, J., Ma, L., Xin, S., Meng, G., Yan, D.M.: Correspondence-free non-rigid point set registration using unsupervised clustering analysis. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. pp. 21199–21208 (2024). <https://doi.org/10.1109/CVPR52733.2024.02003>
26. Zhou, J., Gao, B., Wang, K., Pei, J., Heng, P.A., Qin, J.: Landmark-free preoperative-to-intraoperative registration in laparoscopic liver resection. *IEEE Transactions on Medical Imaging* **44**(11), 4350–4362 (2025). <https://doi.org/10.1109/TMI.2025.3574198>