

SS-IoSR: Self-Supervised Intraoral Scans Repair

Manel Farhat¹ (✉), Achraf Ben-Hamadou¹, Ahmed Rekik¹, Ons Abida¹, and Oussama Smaoui²

¹ Digital Research Center of Sfax Laboratory of Signals, Systems, Artificial Intelligence and Networks Technopark of Sfax, Sfax 3021, Tunisia

`farhat.manel@hotmail.fr`

² Institute for Artificial Intelligence, Biotech Dental Group, Salon de Provence, France

Abstract. Intraoral 3D scans are fundamental in digital dentistry but often contain geometric artifacts such as holes and non-manifold structures due to acquisition constraints and complex dental anatomy. We introduce SS-IoSR, a self-supervised hybrid framework for intraoral scan repair that combines masked autoencoder-based geometric representation learning with a hybrid explicit-implicit reconstruction strategy. The model operates on localized point cloud patches guided by clinician-selected seed points, enabling anatomically consistent reconstruction of missing regions. A differentiable Poisson-based refinement module further improves surface continuity and geometric fidelity. Experimental evaluations on public and in-house intraoral scan datasets show that SS-IoSR consistently outperforms state-of-the-art completion methods across multiple metrics, demonstrating improved reconstruction accuracy and higher structural similarity for clinically relevant dental regions.

Keywords: Intraoral 3D Scans · Mesh reconstruction · Point Cloud Completion · Dental Scan Repair · Digital Dentistry.

1 Introduction

3D intraoral scanning has become central in modern digital dentistry, providing patient-specific, high-resolution 3D reconstruction of dental anatomy for diagnostics, treatment planning, and prosthetic fabrication.

Despite advances in acquisition technology, raw intraoral meshes frequently exhibit imperfections (see Fig. 1(a)), including holes, non-manifold edges, and disconnected components caused by reflective surfaces, limited accessibility, and scanner constraints [8]. Such artifacts often require manual correction or even rescanning to achieve clinically usable models [12].

Unlike standard point cloud benchmarks such as ShapeNet [21] or PCN [27], intraoral scans exhibit highly nonuniform sampling and require precise reconstruction of fine anatomical structures, particularly in interdental and occlusal regions where small geometric errors may affect treatment planning. These characteristics make generic completion methods suboptimal for clinical data. To address these challenges, we propose a self-supervised hybrid local patch-based

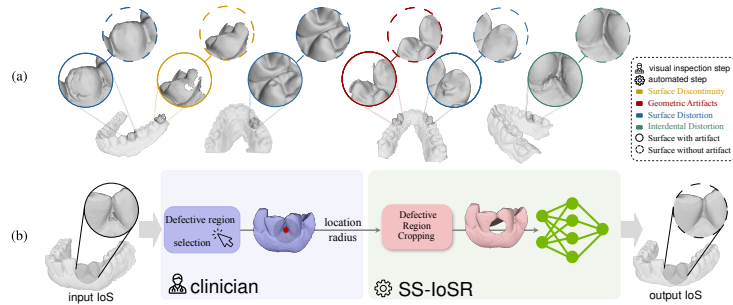


Fig. 1. Graphical abstract. (a): Illustrations of anomalies, each presented with a zoomed-in view highlighting the artifact along with its relevant ground truth. (b): Overview of the proposed repairing pipeline.

completion framework tailored to intraoral geometry. The model is trained using masked reconstruction, enabling learning from unlabeled scans while preserving fine local structures and adapting to density variations.

Contributions: We propose SS-IoSR, a self-supervised framework for repairing 3D intraoral scans in a semi-automatic setting. A clinician initializes the repair process by selecting a seed point corresponding to the approximate centroid of the defect region in a human-in-the-loop manner, ensuring clinical safety and targeted reconstruction while mitigating hallucinations commonly observed in fully automatic approaches. The model then reconstructs anatomically consistent geometry (Fig. 1(b)). To better capture the complementary strengths of geometric reasoning and surface continuity, we introduce a hybrid completion architecture combining explicit geometric modeling for local structural preservation with implicit representations for continuous surface refinement. Finally, to support reproducibility and future research, we release the first paired dataset of intraoral scans before and after repair, providing a benchmark for automatic dental mesh restoration. The project page, including the source code and dataset, is available at: <https://crns-smartvision.github.io/iosr>.

2 Related Work

3D Point Cloud Completion. 3D point cloud completion has been extensively studied in general computer vision, primarily on synthetic object datasets. Early point-based methods such as PCN [27] and TopNet [16], as well as grid-based approaches like GRNet [22], reconstruct shapes using encoder-decoder architectures operating on partial point sets. More recent approaches leverage adversarial learning, transformer architectures, and diffusion priors to improve structural consistency and detail synthesis [10,25,19]. However, these methods are primarily designed for object-level completion under synthetic and relatively

uniform sampling conditions, which limits their robustness when applied to clinically acquired intraoral scans exhibiting localized defects and highly nonuniform geometry.

In the dental domain, geometry generation has also been explored for prosthetic design, such as crown reconstruction methods that generate complete tooth geometries under strong clinical priors and predefined target shapes [23]. However, these approaches address controlled geometric synthesis rather than repairing arbitrary localized defects on existing intraoral surfaces. In contrast, intraoral scan repair focuses on restoring localized anatomical structures on partially corrupted meshes, particularly in clinically critical regions such as interproximal and occlusal surfaces where preserving fine anatomical continuity is essential. Recent dental repair work such as Farhat et al. [6] proposed a framework combining geometric and implicit representations in a supervised global completion setting. However, it does not explicitly exploit self-supervised representation learning or localized patch-based modeling for clinically targeted defect repair. In contrast to global strategies, *SS-IoSR* focuses on localized repair to better preserve high-frequency details that are often oversmoothed by global methods, while learning in a self-supervised manner without relying on labeled data.

Self-Supervised Learning in 3D. Due to the limited availability of annotated clinical data, self-supervised learning (SSL) has become an attractive paradigm for 3D representation learning, inspired by masked modeling in natural language processing [4] and computer vision [9]. Most successful 3D SSL frameworks, including Point-BERT [26] and Point-MAE [13], operate on explicit point-based representations by masking and reconstructing local point patches. Generative variants such as PointGPT [3] and geometry-aware extensions like GeoMAE [17] further build on this explicit patch modeling paradigm. However, these approaches are primarily designed for representation focus on explicit point coordinates reconstruction rather than learning implicit surface functions or performing localized geometric repair.

Implicit 3D Representations. Explicit point-based representations, while flexible, often struggle to enforce surface continuity and consistent topology. In contrast, implicit neural representations, such as DeepSDF [14] and more recently MetricGrids [18], model shapes as continuous volumetric fields by learning functions over spatial coordinates. These methods rely on global class-based shape priors and are limited in capturing fine-grained, case-specific variations in clinical data. Recent implicit parametric models, such as NPHM [7] for head modeling and Zhang et al. [28] for dental reconstruction, leverage strong geometric priors for high-fidelity reconstruction but typically require explicit correspondence, semantic initialization, or part-level supervision.

These methods are particularly effective for holistic 3D reconstruction, producing smooth and topology-agnostic surfaces suitable for high-fidelity modeling. However, most existing implicit reconstruction approaches are formulated for supervised global shape completion, aiming to recover overall object geometry from

partial observations. In contrast, intraoral scan repair targets localized geometric defects on otherwise complete meshes, requiring spatially localized reasoning and tighter integration with explicit surface representations rather than purely global volumetric modeling.

These methods are particularly effective for holistic 3D reconstruction, producing smooth and topology-agnostic surfaces suitable for high-fidelity modeling. However, most existing implicit reconstruction approaches are formulated for supervised global shape completion, aiming to recover overall object geometry from partial observations. In contrast, intraoral scan repair targets localized geometric defects on otherwise complete meshes, requiring spatially localized reasoning and tighter integration with explicit surface representations rather than purely global volumetric modeling.

Taken together, these observations motivate a self-supervised framework that integrates explicit and implicit geometric representations for localized intraoral scan repair. Explicit point-based modeling effectively captures local morphology through patch-level reasoning, whereas implicit representations provide continuous surface modeling at a global scale. We therefore propose *SS-IoSR*, a self-supervised hybrid framework designed for spatially localized and anatomically consistent repair. By learning directly from unlabeled scans, our approach achieves high-precision restoration even under the non-uniform sampling and localized defects characteristic of clinically realistic acquisition conditions.

3 Proposed Approach

As illustrated in Fig. 2, the proposed framework is trained in two phases. In the first phase, we employ a reconstruction-based masked autoencoder (MAE) [9] to reconstruct masked point patches. This pre-training objective encourages the model to learn discriminative local geometric representations, enabling the capture of fine-grained surface details. In the second phase, corresponding to the completion task, the pre-trained MAE serves as the backbone and is integrated with an implicit refinement module. This module enhances the initially reconstructed point cloud by leveraging the continuous modeling capacity of implicit functions to improve geometric fidelity and surface consistency.

3.1 Phase 1: Self-supervised Pre-training

Patch Generation. Most masked point cloud methods [13,29,3] rely on KNN to create fixed-size patches, which struggle with intraoral scans due to their highly non-uniform point densities. To address this, we adopt a fixed-radius patching strategy. This radius is defined jointly with clinical operators to ensure anatomical relevance. Farthest Point Sampling (FPS) on the original point cloud $X \in \mathbb{R}^{W \times 3}$ selects K well-distributed centroids $C \in \mathbb{R}^{K \times 3}$. For each centroid, a fixed-radius neighborhood search is performed, followed by local FPS to subsample a fixed number of points. This process extracts patches $P = \{P_k \in \mathbb{R}^{N \times 3}\}_{k=1}^K$

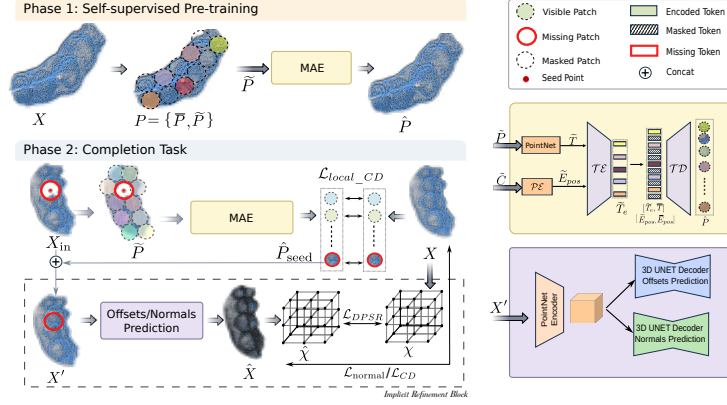


Fig. 2. Overview of the SS-IoSR framework. Phase 1 shows self-supervised MAE for reconstructing input point cloud patches. Phase 2 illustrates the completion task with missing patch reconstruction.

with a consistent geometric extent, ensuring uniform patch cardinality and spatial scale. The extracted patches are then processed patch-by-patch through a masked autoencoder.

Masking and reconstruction. At this stage, a fraction $\gamma \in [0, 1]$ of the patches is randomly masked, resulting in $M = \gamma \cdot K$ masked patches from P . This procedure divides the patches into two subsets: masked patches $\bar{P}$ and visible patches $\tilde{P}$. Only the visible patches $\tilde{P}$ and their centers $\tilde{C}$ are fed into the transformer encoder $\mathcal{TE}$. The visible patch centers are first projected into a D -dimensional embedding space using a positional encoder $\mathcal{PE}$ implemented as an MLP.

$$\tilde{E}_{pos} = \mathcal{PE}(\tilde{C}), \quad \tilde{E}_{pos} \in \mathbb{R}^{(K-M) \times D}. \quad (1)$$

In parallel, the visible patches are encoded using a lightweight PointNet [2], which applies MLP layers followed by a max-pooling operation to generate patch tokens $\tilde{T}$.

To encode the visible tokens $\tilde{T}$ while incorporating spatial information, the positional embeddings $\tilde{E}_{pos}$ are added to the tokens within each Transformer block to preserve spatial structure throughout encoding. The resulting representations are processed by a standard Transformer encoder $\mathcal{TE}$, which models contextual dependencies via self-attention and feed-forward layers. The encoded tokens are given by:

$$\tilde{T}_e = \mathcal{TE}(\tilde{T}, \tilde{E}_{pos}), \quad \tilde{T}_e \in \mathbb{R}^{(K-M) \times D}. \quad (2)$$

After obtaining the encoded visible tokens $\tilde{T}_e$, the decoder reconstructs the masked content. A set of learnable mask token embeddings $\bar{T} \in \mathbb{R}^{M \times D}$ is introduced and concatenated with $\tilde{T}_e$ to form the complete reconstruction sequence.

These mask tokens act as learnable geometric placeholders that serve as spatial queries, enabling the decoder to infer the locations and geometry of missing regions. Their corresponding positional embeddings, $\tilde{E}_{pos}$ and $\bar{E}_{pos}$, are concatenated accordingly to preserve spatial information. The resulting token and positional sequences are then processed by a Transformer decoder $\mathcal{TD}$, defined as:

$$\bar{T}_d = \mathcal{TD}([\tilde{T}_e, \bar{T}], [\tilde{E}_{pos}, \bar{E}_{pos}]). \quad (3)$$

Finally, the decoder output $\bar{T}_d$ is passed through a linear projection head to predict the original point patches $\hat{P}$.

3.2 Phase 2: Completion Task

In the completion task, our model receives as input an incomplete intraoral scan crop X_{in} along with a seed 3D point $\mathcal{S}$, manually selected by a dental professional to indicate the centroid of the missing region. We leverage the pre-trained MAE as an explicit geometry reconstruction module. The encoder processes the visible geometry $\tilde{P}$ and extracts latent feature representations. The fine-tuned decoder reconstructs the masked patches using visible and masked tokens together with positional embeddings derived from the seed point $\mathcal{S}$.

To further refine geometric fidelity, we introduce an implicit refinement module. The incomplete point cloud X_{in} is concatenated with the predicted seed-centered patch $\hat{P}_{seed}$, forming $X' = \text{concat}(X_{in}, \hat{P}_{seed})$. The aggregated point set X' is processed using a lightweight PointNet-based encoder followed by two independent 3D U-Net decoders that separately predict point coordinate offsets and surface normals, producing a refined oriented point cloud $\hat{X}$.

Finally, geometric consistency is further improved by estimating a continuous indicator function grid using Differentiable Poisson Surface Reconstruction (DPSR) [15]. The final 3D mesh is then extracted from this field using the marching cubes algorithm.

3.3 Training Losses

During pre-training, the reconstruction of the masked patches is supervised using a local Chamfer Distance loss $\mathcal{L}_{local_CD}$ between the predicted patches $\hat{P}_{seed}$ and the corresponding ground-truth patches P_{seed} .

The completion task uses four losses: a local Chamfer Distance between predicted and ground-truth patches to reconstruct missing regions, To promote fine-grained geometric detail, we combine a global Chamfer Distance loss $\mathcal{L}_{CD}$ between the predicted dense point cloud $\hat{X}$ and the ground truth X with a Mean Absolute Error loss $\mathcal{L}_{normal}$ on the surface normals, and finally, to enforce geometric consistency, we introduce an implicit representation loss $\mathcal{L}_{DPSR}$, defined as the Mean Squared Error between the predicted indicator grid $\hat{\chi}$ and the Poisson-based ground-truth grid χ . The overall completion task training loss is defined as:

$$\mathcal{L} = \alpha \mathcal{L}_{local_CD} + \beta \mathcal{L}_{CD} + \mathcal{L}_{normal} + \mathcal{L}_{DPSR}. \quad (4)$$

with α and β controlling the contributions of the local and global Chamfer distance losses.

4 Experiments and results

All experiments were conducted on the Teeth3DS dataset [1] and our in-house dataset of 200 intraoral scans with raw and manually cleaned versions for aligner fabrication in orthodontics. For pre-training, we used 30000 spatial crops from complete Teeth3DS scans. For the completion task, 25000 masked crops were generated around seed points with variable defect sizes. All crops were normalized and downsampled to 8000 points using FPS. For both tasks, 40 patches of 200 points were extracted per crop, with patch radius equal to 3 mm, defined jointly with clinical operators to ensure anatomical relevance. The masking ratio during pre-training was 65%. Evaluation was performed on 5000 anomaly-centered crops from our in-house dataset. During training, α was set to 10 and β was gradually increased from 0 to 10 after 50K iterations. The model was trained using Adam ($\text{lr } 5 \times 10^{-4}$, batch size 24) and achieves 0.14s inference time on an NVIDIA RTX 4090 GPU (24GB).

4.1 Evaluation and metrics

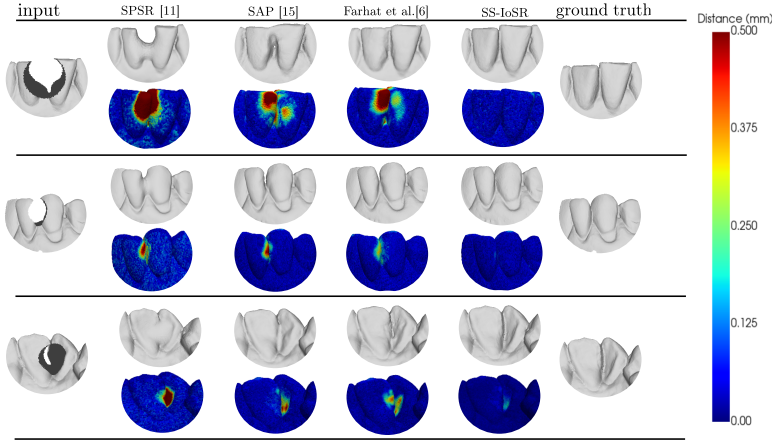
We evaluated our method using standard point cloud completion metrics, including CD-L1 and CD-L2 [5], F-score@10% and F-score@20%, and Density-aware Chamfer Distance (DCD) [20] to account for density variations. As shown in Table 1, we compare our method with several state-of-the-art point cloud completion approaches. The results indicate that SS-IoSR consistently achieves strong performance in terms of distance-based reconstruction errors and F-scores across metrics. Although SAP [15] was not specifically designed for completion tasks, its competitive performance demonstrates the potential of implicit geometric representations for surface modeling. By combining explicit patch-based reasoning with implicit surface refinement, SS-IoSR provides improved geometric reconstruction accuracy and more anatomically consistent results, which is beneficial for intraoral scan completion and downstream dental modeling tasks. As illustrated in Fig. 3, SS-IoSR demonstrates qualitative improvements compared to competing methods. The color-coded distance maps show reduced reconstruction errors, particularly in challenging anatomical regions such as tooth-gingiva boundaries, where preserving fine geometric details is critical. These results further suggest the effectiveness of the proposed framework for accurate intraoral surface reconstruction under clinically realistic conditions. However, SS-IoSR may struggle in cases of extensive structural loss (e.g., a missing crown), where local geometric context becomes insufficient.

4.2 Ablation Study

To evaluate the contribution of each component in SS-IoSR, we conduct an ablation study by progressively removing key modules (see Table 2). Specifically,

Table 1. Comparisons with State-of-the-Art point cloud completion methods.

Methods	$CD-L1$ ($\downarrow$)	$CD-L2$ ($\downarrow$)	DCD ($\downarrow$)	F-score@20% ($\uparrow$)	F-score@10% ($\uparrow$)
TopNet[16]	0.322	0.372	0.580	0.272	0.092
PCN [27]	0.305	0.324	0.541	0.303	0.105
GRNet[22]	0.234	0.172	0.467	0.535	0.143
SAP [15]	0.072	0.036	0.328	0.955	0.641
AdaPoinTr [25]	0.175	0.083	0.410	0.673	0.182
GeoFormer[24]	0.288	0.272	0.516	0.438	0.126
Farhat et al.[6]	0.074	0.032	0.339	0.977	0.78
SS-IoSR	0.062	0.026	0.31	0.983	0.845

**Fig. 3.** Visual results comparing SS-IoSR with state-of-the-art methods (SPSR [11], SAP [15] and Farhat et al.[6])

we compare the full model with variants without self-supervised pre-training (W/O SSL), without the implicit refinement block (W/O IRB), and without the DPSR component (W/O DPSR). Results show that the full SS-IoSR model achieves the best performance. The variant without SSL remains competitive, demonstrating the robustness of the hybrid design, while adding SSL improves feature representation. Removing the implicit refinement block reduces performance due to loss of fine geometric detail, and excluding DPSR causes the largest drop, highlighting the importance of implicit representation for accurate shape modeling.

5 Conclusion

In this work, we introduced SS-IoSR, a self-supervised framework dedicated to repairing intraoral 3D scans using a hybrid design that combines explicit geometric reasoning with implicit surface representations. Through a masked autoencoder pre-training strategy, the model learns geometry-aware and semantically

Table 2. Ablation analysis of SS-IoSR, comparing the complete framework with versions where major components are removed.

Model variants	<i>CD-L1</i> ($\downarrow$)	<i>CD-L2</i> ($\downarrow$)	DCD ($\downarrow$)	F-score@20% ($\uparrow$)	F-score@10% ($\uparrow$)
W/O IRB	0.071	0.052	0.261	0.923	0.578
W/O DPSR	0.143	0.094	0.368	0.783	0.419
W/O SSL	0.068	0.028	0.33	0.979	0.796
Full SS-IoSR	0.062	0.026	0.31	0.983	0.845

rich features in a fully self-supervised manner, enabling effective reconstruction of missing regions in partial intraoral scans. The integration of an implicit refinement module further improves fine-grained surface recovery, producing anatomically consistent and high-fidelity reconstructions. Experimental results demonstrate that SS-IoSR outperforms existing state-of-the-art methods across multiple quantitative metrics, validating the effectiveness of self-supervised learning and the complementarity of explicit and implicit geometric modeling for clinical 3D scan repair.

Disclosure of Interest

The authors declare that they have no competing interests.

References

1. Ben-Hamadou, A., Smaoui, O., Chaabouni-Chouayakh, H., Rekik, A., Pujades, S., Boyer, E., Strippoli, J., Thollot, A., Setbon, H., Trosset, C., et al.: Teeth3ds: a benchmark for teeth segmentation and labeling from intra-oral 3d scans. arXiv e-prints pp. arXiv–2210 (2022)
2. Charles, R.Q., Su, H., Kaichun, M., Guibas, L.J.: PointNet: Deep Learning on Point Sets for 3D Classification and Segmentation. In: 2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 77–85 (Jul 2017)
3. Chen, G., Wang, M., Yang, Y., Yu, K., Yuan, L., Yue, Y.: Pointgpt: auto-regressively generative pre-training from point clouds. In: Proceedings of the 37th International Conference on Neural Information Processing Systems. NIPS ’23, Curran Associates Inc., Red Hook, NY, USA (2023)
4. Devlin, J., Chang, M.W., Lee, K., Toutanova, K.: BERT: Pre-training of deep bidirectional transformers for language understanding. In: Burstein, J., Doran, C., Solorio, T. (eds.) Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers). pp. 4171–4186. Association for Computational Linguistics, Minneapolis, Minnesota (Jun 2019)
5. Fan, H., Su, H., Guibas, L.: A Point Set Generation Network for 3D Object Reconstruction from a Single Image . In: 2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 2463–2471. IEEE Computer Society, Los Alamitos, CA, USA (Jul 2017)

6. Farhat, M., Ben-hamadou, A., Rekik, A., Abida, O., Smaoui, O.: Iosr: End-to-end intraoral scans repairing. In: 36th British Machine Vision Conference 2025, BMVC 2025, Sheffield, UK, November 24-27, 2025. BMVA (2025)
7. Giebenhain, S., Kirschstein, T., Georgopoulos, M., Rünz, M., Agapito, L., Nießner, M.: Learning neural parametric head models. In: Proc. IEEE Conf. on Computer Vision and Pattern Recognition (CVPR) (2023)
8. Gómez-Polo, M., Piedra-Cascón, W., Methani, M.M., Quesada-Olmo, N., Farjas-Abadia, M., Revilla-León, M.: Influence of rescanning mesh holes and stitching procedures on the complete-arch scanning accuracy of an intraoral scanner: An in vitro study. *Journal of Dentistry* **110**, 103690 (2021)
9. He, K., Chen, X., Xie, S., Li, Y., Dollár, P., Girshick, R.: Masked autoencoders are scalable vision learners. In: 2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 15979–15988 (2022)
10. Huang, Z., Yu, Y., Xu, J., Ni, F., Le, X.: PF-Net: Point Fractal Network for 3D Point Cloud Completion. In: 2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 7659–7667 (Jun 2020)
11. Kazhdan, M., Hoppe, H.: Screened poisson surface reconstruction. *ACM Trans. Graph.* **32**(3) (Jul 2013)
12. Narongdej, P., Hassanpour, M., Alterman, N., Rawlins-Buchanan, F., Barjasteh, E.: Advancements in clear aligner fabrication: A comprehensive review of direct-3d printing technologies. *Polymers* **16**(3) (2024)
13. Pang, Y., Wang, W., Tay, F.E.H., Liu, W., Tian, Y., Yuan, L.: Masked autoencoders for point cloud self-supervised learning. In: Computer Vision – ECCV 2022: 23–27, 2022, Proceedings, Part II. p. 604–621 (2022)
14. Park, J.J., Florence, P., Straub, J., Newcombe, R., Lovegrove, S.: DeepSDF: Learning Continuous Signed Distance Functions for Shape Representation. In: 2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 165–174. IEEE, Long Beach, CA, USA (Jun 2019)
15. Peng, S., Jiang, C.M., Liao, Y., Niemeyer, M., Pollefeys, M., Geiger, A.: Shape as points: A differentiable poisson solver. In: Proceedings of the 35th International Conference on Neural Information Processing Systems. pp. 13032–13044. NIPS ’21, Curran Associates Inc., Red Hook, NY, USA (Dec 2021)
16. Tchapmi, L.P., Kosaraju, V., RezaTofighi, H., Reid, I., Savarese, S.: TopNet: Structural Point Cloud Decoder. In: IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 383–392. IEEE, Long Beach, CA, USA (2019)
17. Tian, X., Ran, H., Wang, Y., Zhao, H.: Geomae: Masked geometric target prediction for self-supervised point cloud pre-training. In: Proceedings of the IEEE/CVF (CVPR). pp. 13570–13580 (June 2023)
18. Wang, S., Gao, Y., Li, S., Lv, C., Cai, X., Li, C., Yuan, H., Zhang, J.: Metricgrids: Arbitrary nonlinear approximation with elementary metric grids based implicit neural representation. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 21381–21391 (June 2025)
19. Wei, G., Feng, Y., Ma, L., Wang, C., Zhou, Y., Li, C.: Pcdreamer: Point cloud completion through multi-view diffusion priors. In: 2025 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 27243–27253 (2025)
20. Wu, T., Pan, L., Zhang, J., Wang, T., Liu, Z., Lin, D.: Density-aware chamfer distance as a comprehensive metric for point cloud completion. In: Proceedings of the 35th International Conference on Neural Information Processing Systems. NIPS ’21, Curran Associates Inc., Red Hook, NY, USA (2021)

21. Wu, Z., Song, S., Khosla, A., Yu, F., Zhang, L., Tang, X., Xiao, J.: 3D ShapeNets: A deep representation for volumetric shapes . In: 2015 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 1912–1920. IEEE Computer Society, Los Alamitos, CA, USA (Jun 2015)
22. Xie, H., Yao, H., Zhou, S., Mao, J., Zhang, S., Sun, W.: Grnet: Gridding residual network for dense point cloud completion. p. 365–381. Springer-Verlag, Berlin, Heidelberg (2020)
23. Yang, S., Han, J., Lim, S.H., Yoo, J.Y., Kim, S., Song, D., Kim, S., Kim, J.M., Yi, W.J.: Dcrownformer: Morphology-aware point-to-mesh generation transformer for dental crown prosthesis from 3d scan data of antagonist and preparation teeth. In: MICCAI 2024. pp. 109–119. Springer (2024)
24. Yu, J., Huang, B., Zhang, Y., Li, H., Tang, X., Gao, S.: Geoformer: Learning point cloud completion with tri-plane integrated transformer. In: ACM Multimedia 2024 (2024)
25. Yu, X., Rao, Y., Wang, Z., Lu, J., Zhou, J.: AdaPoinTr: Diverse Point Cloud Completion With Adaptive Geometry-Aware Transformers . IEEE Transactions on Pattern Analysis & Machine Intelligence **45**(12), 14114–14130 (Dec 2023)
26. Yu, X., Tang, L., Rao, Y., Huang, T., Zhou, J., Lu, J.: Point-bert: Pre-training 3d point cloud transformers with masked point modeling. In: 2022 IEEE/CVF (CVPR). pp. 19291–19300 (2022)
27. Yuan, W., Khot, T., Held, D., Mertz, C., Hebert, M.: PCN: Point Completion Network. In: International Conference on 3D Vision (3DV). pp. 728–737 (2018)
28. Zhang, C., Elgharib, M., Fox, G., Gu, M., Theobalt, C., Wang, W.: An implicit parametric morphable dental model. ACM Trans. Graph. **41**(6) (nov 2022)
29. Zhang, R., Guo, Z., Fang, R., Zhao, B., Wang, D., Qiao, Y., Li, H., Gao, P.: Point-m2ae: multi-scale masked autoencoders for hierarchical point cloud pre-training. In: Proceedings of the 36th International Conference on Neural Information Processing Systems. NIPS ’22, Curran Associates Inc., Red Hook, NY, USA (2022)