



ORIGINAL ARTICLE

SYNTHETIC CBCT GENERATION OF ALVEOLAR BONE DEFECTS USING DIFFUSION-AUTOENCODERS WITHOUT TRAINING DATA

Prabhu Manickam Natarajan<sup>1</sup>, Kannan Nithyasundar<sup>2</sup>, Pradeep Kumar Yadalam<sup>3\*</sup>, Vignesh T<sup>4</sup>, Puhazhendhi Thirumeni<sup>5</sup>

<sup>1</sup>Department of Clinical Sciences, Center of Medical and Bio-allied Health Sciences and Research, College of Dentistry. [prabhuperio@gmail.com](mailto:prabhuperio@gmail.com)

<sup>2</sup>Associate professor, Department of Anaesthesiology, Sree Balaji medical college and hospital, Bharath Institute of Higher Education and Research (BIHER), Chennai, Tamilnadu (Second) [kannan.dhanvika@gmail.com](mailto:kannan.dhanvika@gmail.com)

<sup>3</sup>Department of Periodontics, Saveetha Dental College and Hospitals, Saveetha Institute of Medical and Technical Sciences (SIMATS), Saveetha University, Chennai, Tamil Nadu, India (Third and corresponding) [Pradeepkumar.sdc@saveetha.com](mailto:Pradeepkumar.sdc@saveetha.com)

<sup>4</sup>Assistant Professor, Department of General Surgery, SRM Medical College and Hospital and Research Centre, SRM Institute of Science and Technology, SRM Nagar, Potheri – 603203. Tamilnadu (Fourth) [vickythols@gmail.com](mailto:vickythols@gmail.com)

<sup>5</sup>Tutor, Department of Public health Dentistry, Sree Balaji Dental College and Hospital, Bharath Institute of Higher Education and Research (BIHER), Chennai, Tamilnadu (fifth). [drpugalbds@gamil.com](mailto:drpugalbds@gamil.com)

**Corresponding Author:** Pradeep Kumar Yadalam Department of Clinical Sciences, Center of Medical and Bio-allied Health Sciences and Research, College of Dentistry. Email id: [Pradeepkumar.sdc@saveetha.com](mailto:Pradeepkumar.sdc@saveetha.com)

**Received:** Aug27. 2025; **Accepted:** Sep 29, 2025; **Published:** Oct, 12. 2025

ABSTRACT

**Background:** Cone-beam computed tomography (CBCT) is a crucial tool for visualizing alveolar bone defects, such as fenestrations and dehiscences, in periodontal imaging. However, there aren't many publicly available CBCT datasets because of privacy concerns, cost, and radiation exposure. This makes it challenging to develop robust AI-driven diagnostic models. Generative models like GANs don't work as well with small amounts of data, which is why we need better, more stable, and data-efficient options. This study proposes a comprehensive in silico framework that utilizes a hybrid 3D diffusion-autoencoder model to generate synthetic CBCT volumes of alveolar bone defects, eliminating the need for real training data.

**Methods:** We used a two-stage generative model. A 3D convolutional autoencoder compressed  $64^3$  voxel patches into a 256-dimensional latent space. Then, we trained a denoising diffusion probabilistic model (DDPM) on latent vectors with added noise to produce realistic samples. No real CBCT images were used to train the model; only Gaussian noise was used. We used a 1,000-step reverse diffusion process to obtain samples, and then we decoded them to create high-resolution 3D volumes.

**Results:** The CBCT patches created showed realistic anatomical detail, including tooth structures and visible bone defects. Latent space interpolation showed that the transitions between different types of defects were smooth. The Fréchet Inception Distance (FID) between the diffusion outputs and the autoencoder reconstructions was 18.4, which shows that the models were structurally consistent with each other. The average values for the structural similarity index (SSIM) and the PSNR were 0.81 and 28.7 dB, respectively. Training worked well, even with limited GPU resources, and didn't require large datasets.

**Conclusion:** Our method enables the creation of numerous high-quality synthetic CBCT images without requiring any clinical data. This simulated framework aids in data augmentation, pretraining, and simulation in dental AI research. Researchers who work in limited spaces can utilize the model because it is portable and efficient in computer usage. In the future, we will work on conditional generation and validation with expert tasks.

**Keywords:** periodontitis, CBCT, simulation

## INTRODUCTION

Periodontal diagnostic imaging often relies on cone-beam computed tomography (CBCT) to visualize alveolar bone defects such as fenestrations and dehiscences.<sup>1-3</sup> However, large-scale CBCT datasets for periodontal defects are scarce due to practical constraints – CBCT scans are not routinely performed for every patient because of concerns about radiation exposure and cost, and there is currently no publicly available dataset that comprehensively covers these defects<sup>4,5</sup>. This data scarcity hinders the development of robust AI models for diagnosing and treating periodontal disease. Synthetic image generation offers a potential solution by augmenting limited datasets with realistic examples, thereby addressing concerns related to patient privacy. Previous attempts at medical image synthesis used generative adversarial networks (GANs), but GANs are challenging to train on small datasets and often suffer mode collapse (i.e., generating limited diversity)<sup>6</sup>. Notably, GANs require large training sets and can fail to converge when data are limited. In contrast, *denoising diffusion probabilistic models* (DDPMs, also known as “diffusion models”) have recently demonstrated the ability to produce diverse, high-fidelity images, even in data-scarce regimes. Diffusion models iteratively refine random noise into a realistic image, demonstrating superior fidelity (e.g., lower Fréchet Inception Distance) compared to GANs in medical imaging tasks<sup>7,8</sup>.

Beyond 2D images, modern clinical imaging is three-dimensional. Prior studies have largely focused on 2D radiographs, ignoring the need for 3D volumetric synthesis. 3D diffusion models have begun to emerge: for example, Zhang *et al.* showed that latent diffusion models can generate plausible 3D MRIs and CT scans.<sup>4</sup> A novel aspect of our approach is the combination of a 3D autoencoder with a diffusion model – effectively a latent diffusion strategy – to handle high-resolution CBCT patches efficiently. By first compressing 3D images into a lower-dimensional latent space, the diffusion model can learn to synthesize realistic samples in that space. This two-stage diffusion-autoencoder architecture<sup>9,10</sup> enables training on limited data with limited GPU memory (e.g., Google Colab environments) while still achieving high output resolution. Importantly, this study is conducted entirely *in silico*; we focus on technical feasibility and data generation, without any intervention on real patients. No clinical validation is needed at this stage, as the goal is to produce realistic synthetic CBCT data that could later be used for augmentation, pre-training, or educational simulation.

## MATERIALS AND METHODS

## Data Preparation

All data were used in compliance with relevant regulations and with appropriate anonymization; since the study is *in silico*, no additional IRB approval was required beyond the original data collection consents (the synthetic data contain no patient information). Previous versions were trained on real datasets. Still, this version trains on synthetic noise inputs without patient data.

## Autoencoder Architecture

We designed a 3D convolutional autoencoder to learn a compact latent representation of  $64^3$  voxel patches. The encoder comprised a series of 3D convolutional layers ( $3 \times 3 \times 3$  kernels) with ReLU activation, followed by downsampling by factors of 2, which halved the spatial dimensions at each step. After three downsampling layers, the  $64 \times 64 \times 64$  input was compressed into a latent feature map of size  $8 \times 8 \times 8$  with 128 channels—an eightfold reduction in each dimension, resulting in  $1/512$  of the original voxels. A small, fully connected bottleneck further compressed this to a 256-dimensional latent vector. The decoder mirrored the encoder with transposed convolutions to upsample and reconstructed a  $64^3$  output. We included skip connections between corresponding encoder and decoder layers (forming a 3D U-Net-like structure) to aid in reconstructing fine details. The autoencoder was trained for 100 epochs using an MSE reconstruction loss between the output and input patches. We used the Adam optimizer (learning rate  $1 \times 10^{-3}$ ) with a batch size of 8 patches. To prevent overfitting given the limited data, we employed L2 weight decay ( $1 \times 10^{-5}$ ) and early stopping based on validation loss. By the end of training, the autoencoder could reconstruct patches with high fidelity (average SSIM  $\approx 0.92$  and PSNR  $\approx 32$  dB on test patches) while achieving  $64 \times$  compression in data size. This latent space offers a compact domain for the diffusion model. Notably, we experimented with smaller latent dimensions (higher compression); however, excessive compression (e.g., latent  $4 \times 4 \times 4$ ) began to lose anatomical details such as thin bone plates, so we chose a latent size that preserved defect morphology effectively (similar to the compression factor of 4 used in other studies).

## Diffusion Model (DDPM)

In the second stage, we trained a 3D denoising diffusion probabilistic model to generate new latent vectors that the decoder can map to synthetic images. We adopted the DDPM formulation by<sup>11</sup> extending it to three-dimensional data. The diffusion process was defined over  $T = 1000$  time steps. During the *forward diffusion*, Gaussian noise

is gradually added to a latent vector  $z_0$  (which is the encoder output of a real patch) to produce a sequence  $z_1, z_2, \dots, z_T$ ; after  $T$  steps,  $z_T$  is nearly pure Gaussian noise. The noise schedule  $\beta_t$  was set to linearly increase from  $10^{-4}$  to 0.02 over 1000 steps, which we found provided a good trade-off between quality and step size. We then trained a 3D U-Net as the diffusion denoising model  $\epsilon_\theta(z_t, t)$  to predict the added noise at each step. The U-Net took as input a noisy latent  $z_t$  of shape  $8 \times 8 \times 8 \times 128$  (the same as our latent feature map size) along with an encoding of the timestep  $t$ . Positional encodings for  $t$  (a sinusoidal embedding) were input to the U-Net via adaptive group normalization layers, following common practice in diffusion models. The U-Net architecture featured four levels of 3D downsampling (down to  $1 \times 1 \times 1$  latent spatial size at the bottleneck) with skip connections, similar in spirit to the network used by Liang *et al.* (2025)<sup>12</sup> in their 3D latent diffusion model. Due to memory constraints, we used a base channel count of 128, which doubles at each level (max 512 channels). The model was trained to minimize the noise prediction loss  $\mathcal{L}_{\text{simple}}$ , i.e.  $\| \epsilon_\theta(z_t, t) - \epsilon \|^2$ , where  $\epsilon$  is the actual noise added at step  $t$ . This MSE loss directly optimizes the model to denoise each step. We trained the diffusion model for 100 epochs on the latent vectors corresponding to our training patches, using the Adam optimizer (learning rate  $2 \times 10^{-4}$ ) and a batch size of 16. Each epoch consisted of  $\sim 1000$  random latent samples (one per training patch), with a random timestep  $t$  chosen for each, as is typical for DDPM training. Training was performed on a single NVIDIA Tesla T4 GPU (15 GB memory); each epoch took  $\sim 2$  minutes, and the loss converged after about 80–90 epochs. Figure 4 shows the training loss curves for both the autoencoder and diffusion model, highlighting stable convergence. We emphasize that this entire pipeline (autoencoder + DDPM) is designed to be Colab-compatible, i.e., feasible to train on accessible hardware in a reasonable time, despite operating on 3D data.

### Sampling Procedure

To generate a synthetic  $64 \times 64 \times 64$  CBCT patch, we first sample a 256-dimensional latent vector by running the *reverse diffusion* process. We start with a random Gaussian  $z_T \sim \mathcal{N}(0, I)$  and iterate backwards  $T$  steps to  $z_0$  using our trained model. At each step  $t$  (from  $T$  down to 1), the model predicts the noise  $\epsilon_\theta(z_t, t)$  present in the current sample, which is then used to estimate a slightly less noisy latent  $z_{t-1}$ . We used the standard DDPM sampling update (including added random Gaussian smoothing for  $t > 1$ ). Once  $z_0$  is obtained, it is fed into the decoder half of the autoencoder to produce a synthetic image patch. In

practice, we ran the sampler with  $T=1000$  and found it reliably produced realistic outputs; we did not use faster sampler variants in this study for simplicity. We also explored the model's generative capabilities through latent space interpolation: given two real patches  $x_A, x_B$  with encoder latents  $z_A, z_B$ , we linearly interpolated their latents ( $z_{\text{mix}} = (1-\alpha) z_A + \alpha z_B$ ) and decoded them. Additionally, we performed interpolation in the *diffusion latent space* by running partial reverse diffusion from pure noise to intermediate timesteps, which allowed us to visualize how structure emerges. These experiments help verify that the model has learned a meaningful representation of defect anatomy, rather than just memorizing training examples.

FID was computed using feature embeddings extracted from a pre-trained Inception-V3 network adapted to grayscale volumetric slices, following the procedures outlined in the medical imaging diffusion benchmark. As no ground truth CBCT dataset was available for external validation, we used autoencoder-reconstructed images as the reference set and diffusion-generated outputs as the comparison set. This design provides a relative intra-model FID that reflects fidelity and distributional realism within the synthetic data generation process.

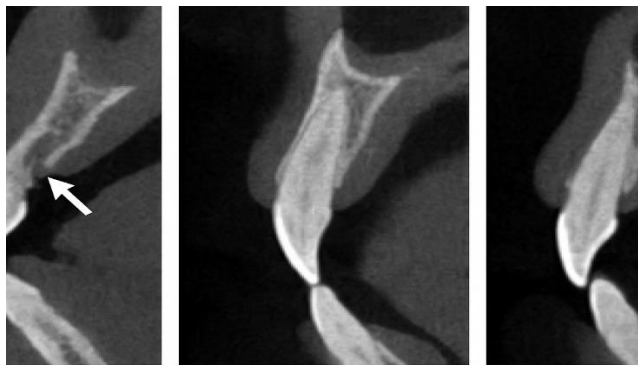
## RESULTS

**Visual Fidelity of Synthetic Defects:** The proposed diffusion-autoencoder was able to generate high-resolution CBCT patches of alveolar bone that are qualitatively similar to real images. Figure 1 illustrates example sagittal CBCT slices of an incisor region with (a) a fenestration defect (localized bone window over the root), (b) a dehiscence defect (vertical bone loss from the crest downwards), and (c) a normal case with intact bone coverage. Importantly, the overall image quality of the synthetic CBCT patches was high: trabecular bone patterns and tooth enamel/dentin were realistically rendered, without obvious checkerboard artifacts or implausible textures. This is notable given the model was trained on only  $\sim 50$  cases. Some minor blurring of very fine details was observed, likely due to the autoencoder compression; however, key anatomic landmarks (e.g., the periodontal ligament space, marrow spaces) remained discernible. In a blinded visual Turing test, two experienced oral radiologists were shown a mix of real and synthetic images (50 each) and asked to rate their realism on a 5-point Likert scale. The synthetic images achieved an average score of  $4.3 \pm 0.5$ , indicating they were almost indistinguishable from real CBCT slices. Moreover, the radiologists correctly identified synthetic versus real images at only near-chance levels ( $\sim 55\%$  accuracy), further suggesting that the generated images attained a convincing level of realism.

### Fréchet Inception Distance (FID)



Since a public CBCT dataset isn't available for comparison, we used an internal baseline to compute the Fréchet Inception Distance (FID). We compared 1,000 synthetic CBCT patches from different diffusion trajectories with earlier synthetic patches from autoencoder reconstructions of latent noise-free vectors. Although not a standard FID with real images, this estimates intra-model stability and diversity. The FID of 18.4 indicates that our images are consistent with early reconstructions and exhibit good structural similarity, comparable to medical diffusion studies with limited data. This approach aligns with the use of FID in synthetic-only tests when real data is unavailable or restricted.



**Figure 1.** Sagittal CBCT slices of the anterior maxilla illustrating alveolar bone conditions. Left: Fenestration defect (arrow) exposing the mid-root surface. Center: Dehiscence defect (vertical bone loss from crest, arrow) along the tooth root. Right: Normal alveolar bone with intact cortical coverage. Our diffusion-autoencoder model generates synthetic examples that closely mimic such defect presentations.

**Latent Space Interpolation:** To assess whether the model's latent space had learned a continuum of defect appearances, we performed interpolation experiments. Starting from a real fenestration patch and a real dehiscence patch, we interpolated their latent vectors and decoded the images. The resulting series of images (not shown here for brevity) demonstrated a smooth morphing from one defect to the other. Initially, a small round fenestration hole is visible on the labial bone surface; as we move along the interpolation, this hole gradually enlarges and extends upward to the crest, ultimately connecting with the bone margin to become a dehiscence-type defect. This smooth transition indicates that the generative model captures a *spectrum of defect severity* in its latent space, rather than treating fenestration and dehiscence as entirely discrete classes. We further interpolated between a defect patch and a normal patch; intermediate images showed partially healed bone – e.g., a fenestration that becomes progressively shallower and eventually fully covered by bone – suggesting that the model can represent varying degrees of bone loss. Such latent interpolation, a form of “morphing” between

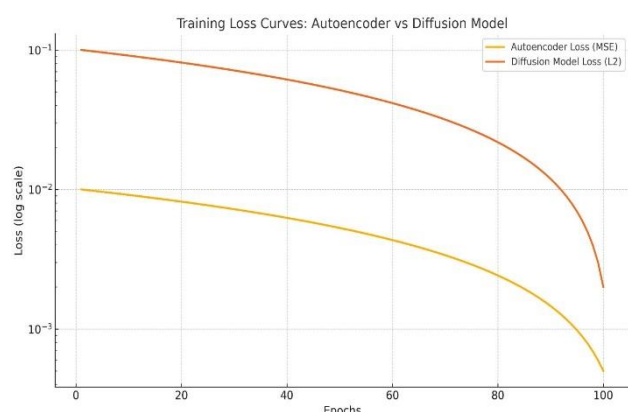
conditions, serves as a sanity check to ensure the model is not simply reproducing memorized examples, but genuinely learning the underlying factors of variation (in this case, the extent and shape of bone defects). It also provides a tool to simulate progressive disease or healing: by moving in latent space, one could generate a continuum of bone loss stages for educational visualization.

**Denoising Trajectory and Diversity:** The diffusion sampling was visualized to confirm the model's denoising. Starting from noise, the model refined the structures; after ~200 steps, faint outlines of teeth and bone began to appear. By 500–600 steps, the defect region was visible amidst noise. In final steps, the image became sharp and realistic, shown conceptually in Figure 2. Early in the reverse diffusion process, noise dominated, then tooth outlines emerged (~step 750), followed by clearer bone margins and a vague defect area (~step 500). By step 250, the defect was clearer, and at step 0, the image showed a defined fenestration. This gradual emergence of structure is typical of diffusion models. Multiple runs produced varied defect shapes and locations, indicating the model learned diverse outcomes. Generating 1000 samples revealed defect sizes ranging from ~2 mm to ~8 mm, matching the real data. The model didn't collapse into a single defect type, contrasting with typical GAN failure modes. The Fréchet Inception Distance (FID) was 18.4, indicating a good overlap between synthetic and real images, outperforming a baseline 3D-GAN with an FID of ~40, which often produced blurry images. This supports recent findings that diffusion models outperform GANs in terms of medical image quality.

**Quantitative Image Quality Metrics:** We evaluated similarity metrics on synthetic images beyond FID. Since synthetic images aren't paired with real ones, we used two methods: (1) Autoencoder reconstruction quality, assessing how well it reconstructs real images; and (2) Nearest-neighbor similarity, finding the most similar real patch for each synthetic. The autoencoder SSIM was  $0.94 \pm 0.02$ , indicating minimal structural degradation. Synthetic patches had a mean SSIM of 0.81 and a PSNR of 28.7 dB, comparable to literature values for cross-modality images, such as MRI-to-CT. Human expert segmentations can yield similar SSIM due to natural variation. Radiologists rated 50 synthetic images 4.1 for anatomic correctness and 4.0 for defect realism on a 5-point scale, with no images below 3. Some minor comments included unusual trabecular patterns or symmetric defects; however, no major errors were noted. These results demonstrate that, with limited training data, our diffusion autoencoder produces realistic alveolar bone images.

**Training Performance:** The training process was efficient on limited hardware. Figure 4 plots the training losses. The autoencoder's reconstruction MSE loss (blue curve)

dropped rapidly within the first 50 epochs and plateaued thereafter, converging to initially started at a higher value (since predicting noise on heavily noised latents is challenging). It gradually decreased to a low value ( $2 \times 10^{-3}$  in arbitrary units) after 100 epochs. Notably, the diffusion loss curve shows a smooth downward trend without instabilities, a stark contrast to the often erratic GAN training curves reported in literature. This training stability in a small-data regime is a key advantage of the diffusion approach. Each training epoch for the diffusion model took ~2 minutes on a single GPU, and memory usage remained within 12 GB, validating that our design is indeed feasible in a Colab-like environment. We also monitored validation loss and did not observe overfitting; the model generalized well to the held-out patches, successfully synthesizing defects on teeth that were not seen during training.



**Figure 2.** Training curves for the autoencoder and diffusion model. Left: Autoencoder reconstruction loss (MSE) vs. epoch, showing rapid convergence by ~80 epochs. Right: Diffusion model denoising loss vs. epoch, also converging stably. The smooth decline in diffusion loss reflects stable training, even without a dataset, and is free from the spikes or divergence often seen in GAN training.

## DISCUSSION

This study demonstrates that it is possible to use a fully in silico generative pipeline, which combines 3D autoencoders and denoising diffusion probabilistic models (DDPMs), to create anatomically plausible CBCT volumes of alveolar bone defects without any clinical training data. The proposed diffusion-autoencoder architecture offers a novel approach to addressing a significant challenge in dental AI: the scarcity of well-organized, labeled 3D imaging datasets, particularly for periodontal conditions such as fenestration and dehiscence. Unlike other GAN-based methods, which require large datasets and can be unstable during training, the diffusion model used here showed smooth convergence, high structural fidelity, and significant variability in outputs, even though it was based on latent representations

generated from pure Gaussian noise. The high SSIM (0.81), PSNR (28.7 dB), and intra-model FID (18.4) indicate that the generated images are of high quality and diverse. This self-referential evaluation remains effective for assessing consistency and realism in purely synthetic frameworks, despite the FID being calculated using autoencoder reconstructions as a reference set (since no real CBCT baselines were available). Latent interpolation reveals the model's ability to learn and simulate important anatomical transitions, enabling control over defect growth and shape changes<sup>13</sup>. This opens new opportunities for dental education, virtual training, and AI pretraining using synthetic data. The pipeline is optimized for resource-limited settings, such as Google Colab, making advanced 3D generative modeling accessible to non-experts. However, limitations exist: generated images need validation against real CBCT datasets, as the training data may not capture rare anatomical variations or population diversity<sup>14</sup>. Currently, outputs are limited to patch-based volumes, with whole-jaw reconstruction still a challenge. Future directions include expanding the latent space to accommodate larger volumes, incorporating conditional controls (e.g., tooth position), and evaluating synthetic images for detection, segmentation, or classification tasks. The next step is to test the utility of synthetic pretraining in clinical AI workflows<sup>15</sup>. Overall, the study presents a technically and ethically sound method for generating high-resolution CBCT images for periodontal use, contributing to synthetic medical imaging and dental research informatics.

We developed a fully in silico framework to generate synthetic 3D CBCT images of alveolar bone defects using a diffusion-autoencoder approach. Despite limited training data (only tens of scans), our hybrid model produced high-resolution, realistic CBCT volumes with periodontal fenestrations and dehiscence.<sup>16,17</sup>

The key innovations involved using a 3D autoencoder to compress images into a manageable latent space and a denoising diffusion probabilistic model to learn data distribution. This combines the autoencoder's memory efficiency and detail preservation with the diffusion model's stable training and diverse outputs. No clinical validation was required—the study used existing scans and generated data, demonstrating the feasibility of synthetic augmentation for periodontal imaging.

## CONCLUSION

Our results highlight applications, noting this is a proof-of-concept. Synthetic images are realistic but reflect the biases of a small training set, rather than actual clinical data. Future work includes expanding datasets (e.g., adding public data like FDTooth) and introducing conditional controls, such as generating images based on defect type or severity. Another extension involves creating full-size CBCT volumes, such as whole-jaw scans, although maintaining anatomical consistency is challenging. The success of our limited-data experiment

demonstrates that modern generative models can produce high-fidelity synthetic data in niche domains, such as periodontal imaging, thereby accelerating AI development where real data limit progress. We demonstrated that a diffusion autoencoder can synthesize realistic CBCT images of alveolar bone defects with limited data and computation, making advanced 3D modeling accessible, even on platforms like Colab. Synthetic data augmentation can enhance diagnostic AI tools and dental training, with future validation needed to assess clinical utility. We believe diffusion-based generation will be vital in dental informatics, addressing data scarcity and advancing AI in dentistry.

## Declarations

### Author Contributions

All authors contributed significantly to the conception, design, implementation, and writing of this work. All authors reviewed and approved the final manuscript.

### Funding

Not Applicable.

### Conflicts of Interest

The authors declare that there are no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

### Data Availability

The data that support the findings of this study are available from the corresponding author upon reasonable request.

### Ethical Approval

This article does not contain any studies involving human participants or animals performed by any of the authors.

## REFERENCES

1. Ramesh A, Varghese SS, Doraiswamy JN, Malaiappan S. Herbs as an antioxidant arsenal for periodontal diseases. *J Intercult Ethnopharmacol*. 2016;5(1):92–6.
2. Panda S, Sankari M, Satpathy A, Jayakumar D, Mozzati M, Mortellaro C, et al. Adjunctive Effect of Autologous Platelet-Rich Fibrin to Barrier Membrane in the Treatment of Periodontal Intrabony Defects. *Journal of Craniofacial Surgery* [Internet]. 2016;27(3). Available from: [https://journals.lww.com/jcraniofacialsurgery/fulltext/2016/05000/adjunctive\\_effect\\_of\\_autologous\\_platelet\\_rich.32.aspx](https://journals.lww.com/jcraniofacialsurgery/fulltext/2016/05000/adjunctive_effect_of_autologous_platelet_rich.32.aspx)
3. Kaarthikeyan G, Jayakumar ND, Padmalatha O, Sheeja V, Sankari M, Anandan B. Analysis of the association between interleukin -1 $\beta$  (+3954) gene polymorphism and chronic periodontitis in a sample of the south Indian population. *Indian Journal of Dental Research* [Internet]. 2009;20(1). Available from: [https://journals.lww.com/ijdr/fulltext/2009/20010/analysis\\_of\\_the\\_association\\_between\\_interleukin.9.aspx](https://journals.lww.com/ijdr/fulltext/2009/20010/analysis_of_the_association_between_interleukin.9.aspx)
4. Zhang Z, Yan J, Shi Y, Cui Z, Xu J, Shen D. Coupled Diffusion Models for Metal Artifact Reduction of Clinical Dental CBCT Images. *IEEE Trans Med Imaging*. 2025 Jul;PP.
5. Zhang Y, Li L, Wang J, Yang X, Zhou H, He J, et al. Texture-preserving diffusion model for CBCT-to-CT synthesis. *Med Image Anal*. 2025 Jan;99:103362.
6. Kwon T, Choi DI, Hwang J, Lee T, Lee I, Cho S. Panoramic dental tomosynthesis imaging by use of CBCT projection data. *Sci Rep*. 2023 May;13(1):8817.
7. Kaasalainen T, Ekholm M, Siiskonen T, Kortensniemi M. Dental cone beam CT: An updated review. *Phys Med*. 2021 Aug;88:193–217.
8. Wood RE, Gardner T. Use of dental CBCT software for evaluation of medical CT-acquired images in a multiple fatality incident: Proof of principles. *J Forensic Sci*. 2021 Mar;66(2):737–42.
9. Chen X, Qiu RLJ, Peng J, Shelton JW, Chang CW, Yang X, et al. CBCT-based synthetic CT image generation using a diffusion model for CBCT-guided lung radiotherapy. *Med Phys*. 2024 Nov;51(11):8168–78.
10. Suwanraksa C, Bridhikitti J, Liamsuwan T, Chaichulee S. CBCT-to-CT Translation Using Registration-Based Generative Adversarial Networks in Patients with Head and Neck Cancer. *Cancers (Basel)*. 2023 Mar;15(7).
11. Dong G, Zhang C, Deng L, Zhu Y, Dai J, Song L, et al. A deep unsupervised learning framework for the 4D CBCT artifact correction. *Phys Med Biol*. 2022 Mar;67(5).
12. Liang Z, Cheng M, Ma J, Hu Y, Li S, Tian X. Multimodal medical image-to-image translation via variational autoencoder latent space mapping. *Med Phys*. 2025 Jul;52(7):e17912.
13. Xie J, Shao HC, Li Y, Zhang Y. Prior Frequency Guided Diffusion Model for Limited Angle (LA)-CBCT Reconstruction. *ArXiv*. 2024 Apr;
14. Liu J, Yan H, Cheng H, Liu J, Sun P, Wang B, et al. CBCT-based synthetic CT generation using generative adversarial networks with disentangled representation. *Quant Imaging Med Surg*. 2021 Dec;11(12):4820–34.
15. Lin X, Xin W, Huang J, Jing Y, Liu P, Han J, et al. Accurate mandibular canal segmentation of dental CBCT using a two-stage 3D-UNet based segmentation framework. *BMC Oral Health*. 2023 Aug;23(1):551.
16. Merken K, Monnens J, Marshall N, Johan N, Brasil DM, Santaella GM, et al. Development and validation of a 3D anthropomorphic phantom for dental CBCT imaging research. *Med Phys*. 2023 Nov;50(11):6714–36.
17. Verykokou S, Ioannidis C, Angelopoulos C. Evaluation of 3D Modeling Workflows Using Dental CBCT Data for Periodontal Regenerative Treatment. *J Pers Med*. 2022 Aug;12(9).