

This is the peer reviewed version of the following article:

Annotating the Inferior Alveolar Canal: the Ultimate Tool / Lumetti, Luca; Pipoli, Vittorio; Bolelli, Federico; Grana, Costantino. - 14233:(2023), pp. 525-536. (Intervento presentato al convegno 22nd International Conference on Image Analysis and Processing (ICIAP 2023) tenutosi a Udine, Italy nel Sep 11-15) [10.1007/978-3-031-43148-7_44].

Springer

Terms of use:

The terms and conditions for the reuse of this version of the manuscript are specified in the publishing policy. For all terms of use and more information see the publisher's website.

15/09/2025 21:27

(Article begins on next page)

Annotating the Inferior Alveolar Canal: the Ultimate Tool

Luca Lumetti¹, Vittorio Pipoli^{1,2}, Federico Bolelli¹, and Costantino Grana¹

¹ University of Modena and Reggio Emilia, Italy

`{name.surname}@unimore.it`

² University of Pisa, Italy

`{name.surname}@phd.unipi.it`

Abstract. The Inferior Alveolar Nerve (IAN) is of main interest in the maxillofacial field, as an accurate localization of such nerve reduces the risks of injury during surgical procedures. Although recent literature has focused on developing novel deep learning techniques to produce accurate segmentation masks of the canal containing the IAN, there are still strong limitations due to the scarce amount of publicly available 3D maxillofacial datasets. In this paper, we present an improved version of a previously released tool, IACAT (Inferior Alveolar Canal Annotation Tool), today used by medical experts to produce 3D ground truth annotation. In addition, we release a new dataset, *ToothFairy*, which is part of the homonymous MICCAI2023 challenge hosted by the Grand-Challenge platform, as an extension of the previously released *Maxillo* dataset, which was the only publicly available. With ToothFairy, the number of annotations has been increased as well as the quality of existing data.

Keywords: IAC · CBCT · 3D Segmentation · Annotation Tool

1 Introduction

The placement of dental implants within the jawbone is a common surgical procedure that can raise different complications due to the presence of the Inferior Alveolar Nerve (IAN). Such nerve lies close to molars roots, necessitating meticulous preoperative planning to avoid causing any type of damage. Hence, an accurate segmentation of the bone cavity containing the IAN is crucial to avoid nerve injuries during surgery [8, 13]. Such a cavity is identified as the Inferior Alveolar Canal (IAC).

The IAC segmentation, usually performed by radiologic technologists, is obtained from a 3D Cone Beam Computed Tomography (CBCT) scan by manually drawing a line on a 2D projection of the original volume. This type of annotation is referred to as *sparse* or 2D annotation (Fig. 1a) and provides medical experts with an *approximate localization* of the IAN’s position and its distance from the molars. Although 3D annotations (Fig. 1b) would provide exact knowledge about IAC shape and position, enabling meticulous surgical plan, they are often unavailable due to the burden of time and effort required to obtain them.

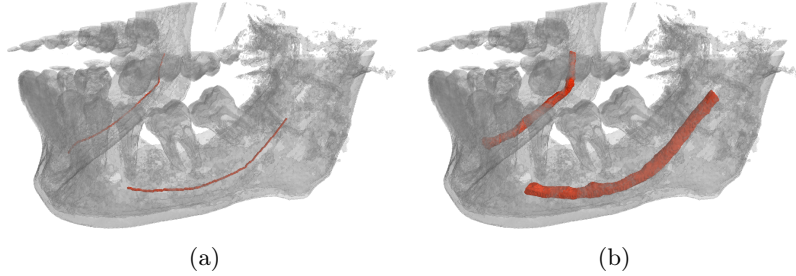


Fig. 1. An example of sparse (a) and dense (b) ground truth annotations of the same volume. Both are obtained from (different) 2D views of the data and later re-projected to the 3D volume. More specifically, the sparse annotation is extracted from a single panoramic view (Fig. 2c), while the dense annotation is obtained from multiple cross sectional views (Fig. 3).

Therefore, the automatic segmentation of the IAC represents an active research field, as it holds the potential to revolutionize the surgical planning process, supporting the maxillofacial daily practice.

Although Convolutional Neural Networks (CNNs) have provided amazing results for both 2D and 3D segmentation, alongside several more computer vision and healthcare tasks [2, 3, 4, 17, 18, 19, 25, 26, 27], developing Deep Neural Networks (DNNs) for the automatic segmentation of the IAC is a challenging task due to the lack of publicly available 3D annotated data³. As the reader may know, collecting and labelling data is a time-consuming and resource-intensive process. Moreover, publicly releasing medical-related information raises privacy issues. To advance the research in IAC automatic segmentation and improve patient surgical outcomes, it is crucial to develop tools and applications that can support the creation of such high-quality datasets.

This paper addresses this issue by enhancing an existing software used by medical experts to produce IAC annotations. We demonstrate that the proposed additional features, later detailed in Sec. 3, allow us to both significantly reduce the annotation time and improve the overall quality of the resulting segmentation. Specifically, radiologists are now able to identify previously undetectable canal sections with an average increase in annotated voxels of 61.9%. Among others, one of the main features introduced in the software is the automatic segmentation of the canal from which technicians can start a simplified annotation process. More specifically, we integrate an improved version of the segmentation approach originally proposed in [7].

It is worth mentioning that the proposed tool, IACAT 2.0, led to the generation of a 3D-segmented-CBCTs dataset, *ToothFairy*, which is part of the homonymous MICCAI2023 challenge. Training state-of-the-art segmentation models with

³ At the moment of writing this paper, there is only one single dataset publicly available: <https://ditto.ing.unimore.it/maxillo/>

such annotations led to significant performance boost both on Dice and Intersection over Union (IoU) scores.

2 Related Work

Previous studies have proposed different architectures for the automatic segmentation of the inferior alveolar canal [7, 9, 14, 15, 16].

In [7] and [9], a three steps training procedure is proposed. During the first step, identified as *deep label expansion*, the network is fed with CBCT volumes paired with the corresponding sparse 2D labels and trained to generate 3D dense annotations. During the second step, the 3D synthetic labels are employed to perform a pre-training of the segmentation network that is finally fine-tuned with the 3D annotations performed by medical experts.

On the other hand, Kwak *et al.* [15], Jaskari *et al.* [14], and Lahoud *et al.* [16] proposed more straightforward approaches. More specifically, [15] simply compares different types of existing 2D and 3D architectures (*e.g.*, SegNet [1] and U-Net [28]) using a private dataset. Instead, a standard 3D-UNet [6] model with some slight improvements tailored for the specific task has been adopted by both [14] and [16].

A significant challenge in the field is represented by the lack of publicly available data and source code. Indeed, most of the aforementioned papers do not provide neither the annotations nor the information about how they have been obtained, which hinders the ability to reproduce their experiments and evaluate the effectiveness of the proposed models. The absence of open-source approaches for IAC segmentation is a major obstacle that must be addressed to facilitate progress in this area of research.

To the extent of annotating 3D data, previous works have relied on the use of proprietary software such as Photoshop⁴ and Invivo⁵ [23, 30], which can be tedious, time-consuming, and not tailored for the specific task. Moreover, even when they propose a novel methodology to annotate such data, they do not release the source code of their implementation [14, 15, 16].

To address this issue, we present a new annotation tool specifically designed for the IAN canal. The tool provides the user with the capability of processing and visualize CBCT data exported in DICOM format and guides him toward the annotation of axial images, panoramic views, and cross-sectional images. The annotated data can be easily exported to be employed in different tasks, including the training of deep learning models.

3 The Ultimate Annotation Tool

With this work, we present IACAT 2.0, an improved version of the tool described in [20]. The proposed features are detailed in the following of this Section, highlighting the rationales behind the design choices and the benefit they introduce.

⁴ <https://www.photoshop.com>

⁵ <https://www.anatomage.com/invivo>

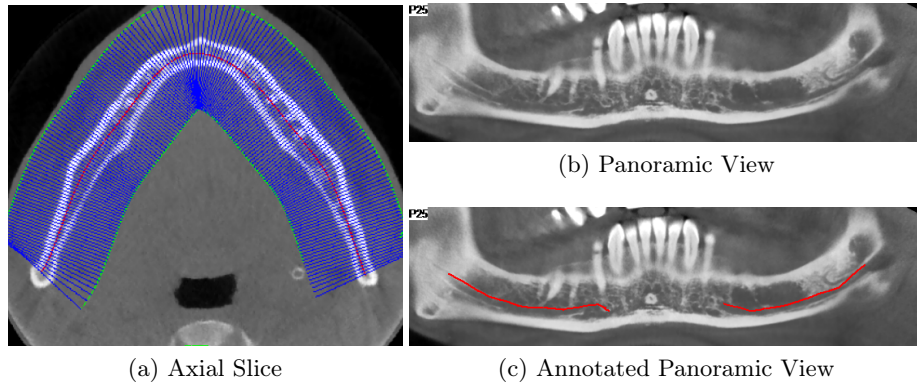


Fig. 2. 2D annotation of the IAC. (a) depicts an axial slice of the CBCT volume. The red curve, called panoramic base curve, identifies the jawbone. (b) is the panoramic view obtained from the CT-volume displaying voxels of the curved plane generated by the base curve and orthogonal to the axial view. (c) is the same view as (b) showing a manual annotation of the IAN performed by an expert technician.

3.1 Preliminaries

To better introduce our proposals, we summarize the entire annotation flow:

1. After loading the input data, the arch approximation that better describes the canal course is identified inside one of the axial planes constituting the volume. The output is a one-pixel thick curve crossing the dental arch which is approximated with a polynomial (Fig. 2a in red). The curve is automatically generated by the tool and manually adjusted only when needed.
2. Sampling the polynomial, the tool thus generates a Catmull-Rom spline. For each point of the spline, a perpendicular line on the axial plane is computed (Fig. 2a in blue). These lines are identified as *Cross-Sectional Lines* (CSLs).
3. CSLs are used for *Multi Planar Reformations* (MPRs) to generate *Cross-Sectional Views* (CSVs). These views are 2D images obtained interpolating the values of the respective base lines (CSL) across the whole volume height. An example of CSV is provided in Fig. 3 (bottom-right), it corresponds to the plane identified in green on the panoramic view (top-right, same Figure).
4. For each CSV, a closed Catmull-Rom spline is finally drawn to annotate the position of the IAC (green closed lines of Fig. 3).
5. The splines are saved as the coordinates of their control points. The final smooth and precise ground-truth volume constituting the dataset is generated from this set of points by means of the α -shape algorithm [10].

Our contributions to the annotation pipeline can be summarized as follows. We integrated an automatic prediction of the IAC based on PosPadUNet3D [7], a state-of-the-art deep learning model for the task. Moreover, to improve the segmentation results, we introduced enhancements in the whole automatic annota-

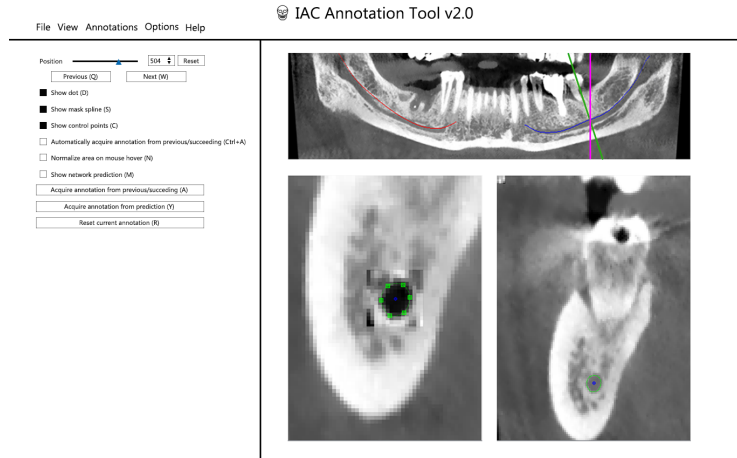


Fig. 3. A view of the IACAT 2.0. The panoramic view (top-right) roughly identifies the canal position (left branch in red, right branch in blue). On the same view, purple and green lines identify the position of a straight and tilted (to be orthogonal to the canal slope) CSVs, respectively. The tilted CSV corresponding to the green line, visualized in the lower part of the screen with two different zoom levels, is intended to produce the annotation by drawing a closed spline. The left-most part of the windows contains control buttons to perform different actions.

tion pipeline. This way, technicians are able to acquire annotations from model predictions, significantly reducing the annotation efforts to a mere adjustment. Additional mechanisms — *e.g.*, zoom in/out and local contrast-stretching— have also been introduced to improve eye-driven identification.

3.2 Automatic Segmentation of the IAC

Compared with the pipeline of PosPadUNet3D, the segmentation of the IAC has been improved through the implementation of ad-hoc pre- and post-processing techniques. Such additions are here introduced and detailed in the following Sections. The pre-processing technique, known as *Distance Transform*, mitigates the sparsity of the 2D annotation fed into the model during the *deep label propagation* step. On the other hand, the post-processing technique, called *Hann Windowing* [24], tackles the artifacts caused by patch-based learning. It achieves this by multiplying the frames with a window function. Together, these techniques have resulted in a more accurate and efficient automatic segmentation of the IAC.

Before digging into the details of the proposed improvements, it is worth noticing that the automatic segmentation of the IAN intervenes at point 4 of the annotation flow presented in Sec. 3.1. More specifically, technicians can visualize the prediction as depicted with red in Fig. 4 and chose whether to generate the closed spline automatically from it, or only use the annotation as a reference.

Positional UNet with Padding. Unlike traditional UNet-inspired models [5, 6, 11, 12, 22, 28, 29], PosPadUNet3D incorporates patch positional information in the bottleneck and employs padded convolution to preserve tensor dimensionality.

As introduced in Sec. 2, the training procedure articulates in three main stages. Initially, the model is fed with a concatenation of the CBCT patches and the corresponding 2D annotations to obtain a dense 3D segmentation, this process is identified as *deep label expansion*. The obtained segmentation is used as the ground truth during the next training step, enabling the use of weakly annotated volumes, such as scans with only 2D annotations. Finally, in the third step, the model is fine-tuned using scans that have 3D labels provided by technicians. When compared to existing state-of-the-art models, PosPadUNet3D demonstrate better performance in segmenting the IAC. Therefore it represents as an efficient and reliable method to implement the *acquire annotation from prediction* in IACAT 2.0.

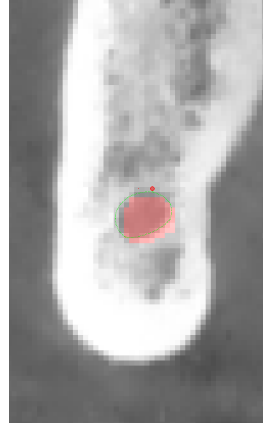


Fig. 4. CSV depicting the annotation performed by medical experts (in green), and the automatic prediction of the network (in red). The dark-red dot is the 2D sparse annotation.

Distance Transform. Unfortunately, the sparse annotations employed for deep label expansion are scrimpy: more than 99.9% of the volume is annotated as background, thus it is not well suited for standard convolutional neural networks nor skip connections. Moreover, during patch-based training, the probability that a patch uniformly sampled from the original volume contains a relevant part of the annotation is low, and there will be mostly empty patches with no information at all, providing just useless computations in the pipeline.

To spread the information contained in the sparse annotation, we propose to add as pre-processing step the Euclidean distance transform. Such transformation takes a binary N-dimensional volume as input and, for each background element, efficiently computes its distance to the closest foreground label. This addition allows the network to gather critical information about the IAC position even in the early stage of the network, and also when dealing with patches that do not directly contain sparse labels. Such an approach enables faster convergence during training and improves final results. The distance transform is defined as:

$$d(x_i) = \min_{x_j} f(x_i, x_j), \quad x_i \in \mathcal{M}_{\text{bg}}, \quad x_j \in \mathcal{M}_{\text{fg}} \quad (1)$$

where x_i refers to all the pixels of the background $\mathcal{M}_{\text{bg}}$ and x_j refers to all the pixels of the foreground $\mathcal{M}_{\text{fg}}$. The function f is the chosen distance function, which in our case corresponds to the standard Euclidean distance.

Hann Windows. Another tackled problem regards the artifacts produced during the patch-based training. When the threshold is removed from the output of the network, it is possible to better notice different inaccuracies produced by the network near the borders of each patch. Additionally, when all the patches are aggregated together, it is possible to notice that the distribution of errors lies exactly where the canal must be predicted close to the borders.

A similar, unrelated, problem also happens in audio encoding, where sub-frames of an audio track are encoded separately and merged afterwards. This procedure causes non-zero values to appear near sub-frame boundaries and is called *spectral leakage*.

As done for audio encoding, we propose to adopt the Hann window function to overcome such a limitation. The Hann window function is defined as follows:

$$W_{\text{Hann}}(i) = \frac{1}{2} \left(1 - \cos \frac{2\pi i}{I} \right)$$

where i is an element in the interval I . The function is symmetric, with a maximum value of 1 in the middle of the window and a minimum value of 0 at the edges. Another interesting property is that the sum of two Hann windows shifted by $\frac{I}{2}$ (50%) is equal to a rectangular window of width I and height 1:

$$W_{\text{Hann}}(i) + W_{\text{Hann}}\left(i + \frac{I}{2}\right) = 1$$

Such a property is employed in audio encoding to remove the border artifacts by simply multiplying overlapped (by 50%) frames with the Hann function before summing them. This approach, which is defined in 1D for the audio, can be extended to be multi-dimensional and thus applied to 3D images:

$$W_{\text{Hann}}(i, j, k) = W_{\text{Hann}}(i)W_{\text{Hann}}(j)W_{\text{Hann}}(k)$$

where i, j, k identify a point in the space. To avoid numerical issues, the windows function applied in the image space must deal with border cases, ensuring that the sum is always one.

3.3 Localized Contrast Stretching and Zoom

The precise segmentation of the IAC in real-life scenarios can be challenging, even for experienced domain experts. This is mainly due to the low contrast in the area where the IAC is located. It is not rare to encounter CSV where, although present, the alveolar canal turns out to be indistinguishable or hard to be localized precisely (see Fig. 5b as an example). This phenomenon typically occurs because bone density is higher in specific locations, due to ageing or other patient-related conditions. However, the information underlying these regions can sometimes be revealed through the application of local contrast stretching. Therefore, we integrated into IACAT 2.0 the ability to increase the contrast range in a given area by applying the following formula:

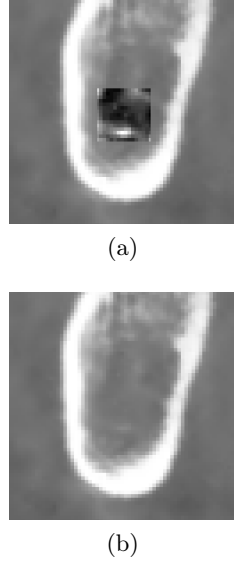


Fig. 5. CSVs with and without localized contrast stretching. (a) depicts a local contrast stretching function applied to the image (b).

region, the *Zoom* function facilitates the inspection of hard-to-annotate regions of the volumes.

The integration of the aforementioned functionalities provides the annotator with powerful tools to overcome the main challenges posed by difficult-to-detect IAC regions.

4 Evaluating IACAT 2.0

A team of 5 maxillofacial experts with more than 5 years of experience have been engaged to perform annotations using both the old and new version of the tool on 40 CBCTs of the public Maxillo dataset. The comparison underlines that annotations obtained with IACAT v1.0 suffered from several issues, including disconnected components and under-annotations that often occur near the terminal parts of the canal (see Fig. 6 as a reference). To produce a quantitative comparison, we introduce the following metric:

$$\text{Increase \%} = \frac{V_{i,\text{New}} - (V_{i,\text{Old}} \cap V_{i,\text{New}})}{V_{i,\text{Old}}} \cdot 100 \quad (5)$$

where $V_{i,\text{New}}$ and $V_{i,\text{Old}}$ are respectively the ground truth annotations of patient i in the corresponding dataset. This gives us a measure of how many more voxels

$$v_{\max} = \max_{i=x-w/2}^{x+w/2} \max_{j=y-h/2}^{y+h/2} I_{i,j} \quad (2)$$

$$v_{\min} = \min_{i=x-w/2}^{x+w/2} \min_{j=y-h/2}^{y+h/2} I_{i,j} \quad (3)$$

where $I \in \mathcal{M}^{M \times N}$ is the matrix representing the original image, w and h are the width and height of the selected window, and (x, y) is the position of the window center. Given that, for each pixel $I_{i,j}$ of the window, the contrast stretching formula applied is:

$$I'_{i,j} = \frac{I_{i,j} - v_{\min}}{v_{\max} - v_{\min}} \quad (4)$$

where $I'_{i,j}$ is the contrast-stretched pixel value at position (i, j) . The window height and width are independently modifiable, to let the user fit the windows exclusively over the region of interest.

As the IAC might be small in some cases, a zoom functionality is also introduced to increase the precision of the carried-out annotation.

While the *contrast stretching* function allows the annotator to enhance the contrast of a specific re-

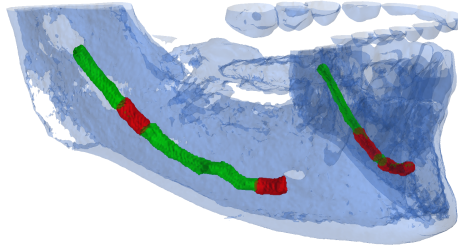


Fig. 6. Comparison of the annotation obtained using the two versions of the tool on volume P95 of the publicly available Maxillo dataset. In red the annotation obtained with the old tool, in green the voxels of the annotation that have been added thanks to the proposed improvements.

have been annotated w.r.t. the annotations performed with the older version of the tool. On average, 61.9% more voxels per volume are selected as being part of the canal. In Fig. 6 an example of the aforementioned discrepancy is depicted.

Additionally, the time spent annotating each one of the 40 CBCTs has been recorded. The average time of 22 minutes per volume required when using IACAT v1.0 has been lowered to around 8 minutes per volume with IACAT 2.0, highlighting once again the benefits of the features introduced.

4.1 Inter-Agreements

To understand the room for improvement of any novel deep learning model, a human-baseline score must be defined. Such a score has been created by using two annotations of the same volume produced by different medical experts, and then computing the Dice and IoU scores among them. Indeed, we produced two annotations for 6 patients and obtained an average IoU of 0.70 and an average Dice of 0.81.

4.2 *ToothFairy* - A New Dataset

As previously stated, an additional contribution of this paper is the generation of a new dataset, *ToothFairy*, obtained by means of IACAT 2.0 using the data already released with the Maxillo dataset. *ToothFairy* counts 153 3D densely annotated CBCTs, *i.e.*, 62 more w.r.t.

Table 1. Performance comparison of different models trained on Maxillo or *ToothFairy* dataset.

Model	Maxillo		ToothFairy	
	IoU	Dice	IoU	Dice
AttentionUNet [22]	0.576	0.731	0.612	0.759
UNet++ [31]	0.542	0.703	0.550	0.710
UNet [28]	0.635	0.777	0.643	0.783
VNet [21]	0.524	0.688	0.558	0.716
PosPadUNet3D [7]	0.652	0.789	0.663	0.797

the original dataset⁶. Regarding the 91 volumes in common with the Maxillo dataset, it is worth noting that 40 of them underwent re-segmentation using IACAT 2.0, while the other 51 annotations remained unchanged. With the aim of demonstrating the value of the new dataset, Tab. 1 compares the performance of different publicly-available state-of-the-art segmentation models trained with both the old and the new datasets. Results demonstrate that the use of the newly produced data is beneficial also for the training of deep neural networks. Notably, the improvement of PosPadUNet3D [7] is relatively modest compared to the other evaluated techniques. This can be explained by considering that its performance are close to the previously defined inter-agreement score (Sec. 4.1). Consequently, considering the human baseline which approximately corresponds to a Dice score of 0.81, the advancement achieved by PosPadUNet3D from 0.79 to 0.80 retains its significance.

5 Conclusion

In this work, we have introduced IACAT 2.0, an innovative tool for the annotation of the inferior alveolar canal in CBCT scans, which enhances the quality and expedites the annotation process. We have also presented ToothFairy, a new dataset of IAC 3D segmentation, which improves the quantity and quality of publicly available annotated CBCT scans. The proposed tool incorporates several novel functionalities, such as *acquire from prediction*, which uses the prediction of a state-of-the-art model to assist the annotation process, and *localized contrast stretching*, which enhances the contrast of dark regions to reveal hidden parts of the alveolar canal, simplifying the annotation process.

The carried out evaluation revealed that, by using IACAT 2.0, medical experts are able to identify previously undetectable canal regions. On average, the new annotations counts 61.9% more voxels and requires $\frac{1}{3}$ of time to be generated. Finally, to highlight the benefits introduced by our tool, we compared state-of-the-art segmentation models trained with and w/o ToothFairy dataset. Significant performance boost have been achieved by all the models on both Dice and IoU scores when using IACAT 2.0-generated data.

Acknowledgements This project has received funding from the Department of Engineering “Enzo Ferrari” of the University of Modena through the FARD-2022 (Fondo di Ateneo per la Ricerca 2022).

References

1. Badrinarayanan, V., Kendall, A., Cipolla, R.: SegNet: A Deep Convolutional Encoder-Decoder Architecture for Image Segmentation. *IEEE Transactions on Pattern Analysis and Machine Intelligence* **39**(12), 2481–2495 (2017) 3

⁶ Please note that the Maxillo dataset released a total of 343 CBCTs proving a 3D annotation only for 91 of them.

2. Barraco, M., Stefanini, M., Cornia, M., Cascianelli, S., Baraldi, L., Cucchiara, R.: CaMEL: Mean Teacher Learning for Image Captioning. In: Proceedings of the International Conference on Pattern Recognition (2022) [2](#)
3. Bolelli, F., Baraldi, L., Grana, C.: A Hierarchical Quasi-Recurrent approach to Video Captioning. In: 2018 IEEE International Conference on Image Processing, Applications and Systems (IPAS). pp. 162–167. IEEE (Dec 2018) [2](#)
4. Bontempo, G., Porrello, A., Bolelli, F., Calderara, S., Ficarra, E.: DAS-MIL: Distilling Across Scales for MIL Classification of Histological WSIs. In: Medical Image Computing and Computer Assisted Intervention – MICCAI 2023 (2023) [2](#)
5. Chen, J., Lu, Y., Yu, Q., Luo, X., Adeli, E., Wang, Y., Lu, L., Yuille, A.L., Zhou, Y.: TransUNet: Transformers Make Strong Encoders for Medical Image Segmentation. arXiv preprint arXiv:2102.04306 (2021) [6](#)
6. Çiçek, Ö., Abdulkadir, A., Lienkamp, S.S., Brox, T., Ronneberger, O.: 3D U-Net: Learning Dense Volumetric Segmentation from Sparse Annotation. In: Medical Image Computing and Computer-Assisted Intervention – MICCAI 2016. pp. 424–432. Springer (2016) [3](#), [6](#)
7. Cipriano, M., Allegretti, S., Bolelli, F., Pollastri, F., Grana, C.: Improving Segmentation of the Inferior Alveolar Nerve through Deep Label Propagation. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 21137–21146. IEEE (2022) [2](#), [3](#), [4](#), [9](#), [10](#)
8. Crowson, M.G., Ranisau, J., Eskander, A., Babier, A., Xu, B., Kahmke, R.R., Chen, J.M., Chan, T.C.: A contemporary review of machine learning in otolaryngology–head and neck surgery. *The Laryngoscope* **130**(1), 45–51 (2020) [1](#)
9. Di Bartolomeo, M., Pellacani, A., Bolelli, F., Cipriano, M., Lumetti, L., Negrello, S., Allegretti, S., Minafra, P., Pollastri, F., Nocini, R., et al.: Inferior Alveolar Canal Automatic Detection with Deep Learning CNNs on CBCTs: Development of a Novel Model and Release of Open-Source Dataset and Algorithm. *Applied Sciences* **13**(5), 3271 (2023) [3](#)
10. Edelsbrunner, H., Kirkpatrick, D., Seidel, R.: On the shape of a set of points in the plane. *IEEE Transactions on Information Theory* **29**(4), 551–559 (1983) [4](#)
11. Guan, S., Khan, A.A., Sikdar, S., Chitnis, P.V.: Fully Dense UNet for 2-D Sparse Photoacoustic Tomography Artifact Removal. *IEEE Journal of Biomedical and Health Informatics* **24**(2), 568–576 (2019) [6](#)
12. Hatamizadeh, A., Nath, V., Tang, Y., Yang, D., Roth, H.R., Xu, D.: Swin UNETR: Swin Transformers for Semantic Segmentation of Brain Tumors in MRI Images. In: International MICCAI Brainlesion Workshop. pp. 272–284. Springer (2022) [6](#)
13. Hwang, J.J., Jung, Y.H., Cho, B.H., Heo, M.S.: An overview of deep learning in the field of dentistry. *Imaging Science in Dentistry* **49**(1), 1–7 (2019) [1](#)
14. Jaskari, J., Sahlsten, J., Järnstedt, J., Mehtonen, H., Karhu, K., Sundqvist, O., Hietanen, A., Varjonen, V., Mattila, V., Kaski, K.: Deep Learning Method for Mandibular Canal Segmentation in Dental Cone Beam Computed Tomography Volumes. *Scientific Reports* **10**(1), 5842 (2020) [3](#)
15. Kwak, G.H., Kwak, E.J., Song, J.M., Park, H.R., Jung, Y.H., Cho, B.H., Hui, P., Hwang, J.J.: Automatic mandibular canal detection using a deep convolutional neural network. *Scientific Reports* **10**(1), 5711 (Mar 2020) [3](#)
16. Lahoud, P., Diels, S., Niclaes, L., Van Aelst, S., Willems, H., Van Gerven, A., Quirynen, M., Jacobs, R.: Development and validation of a novel artificial intelligence driven tool for accurate mandibular canal segmentation on CBCT. *Journal of Dentistry* **116**, 103891 (2022) [3](#)

17. Lovino, M., Ciaburri, M.S., Urgese, G., Di Cataldo, S., Ficarra, E.: DEEPrior: a deep learning tool for the prioritization of gene fusions. *Bioinformatics* **36**(10), 3248–3250 (2020) [2](#)
18. Lovino, M., Urgese, G., Macii, E., Di Cataldo, S., Ficarra, E.: A Deep Learning Approach to the Screening of Oncogenic Gene Fusions in Humans. *International journal of molecular sciences* **20**(7), 1645 (2019) [2](#)
19. Marconato, E., Bontempo, G., Ficarra, E., Calderara, S., Passerini, A., Teso, S.: Neuro Symbolic Continual Learning: Knowledge, Reasoning Shortcuts and Concept Rehearsal. In: *International Conference on Machine Learning (ICML)* (2023) [2](#)
20. Mercadante, C., Cipriano, M., Bolelli, F., Pollastri, F., Anesi, A., Grana, C.: A Cone Beam Computed Tomography Annotation Tool for Automatic Detection of the Inferior Alveolar Nerve Canal. In: *16th International Conference on Computer Vision Theory and Applications-VISAPP 2021*. vol. 4, pp. 724–731. *SciTePress* (2021) [3](#)
21. Milletari, F., Navab, N., Ahmadi, S.A.: V-Net: Fully Convolutional Neural Networks for Volumetric Medical Image Segmentation. In: *2016 Fourth International Conference on 3D Vision (3DV)*. pp. 565–571. *IEEE* (2016) [9](#)
22. Oktay, O., Schlemper, J., Le Folgoc, L., Lee, M., Heinrich, M., Misawa, K., Mori, K., McDonagh, S., Hammerla, N.Y., Kainz, B., Glocker, B., Rueckert, D.: Attention U-Net: Learning Where to Look for the Pancreas. In: *Medical Imaging with Deep Learning* (2022) [6](#), [9](#)
23. Park, J.S., Chung, M.S., Hwang, S.B., Lee, Y.S., Har, D.H.: Technical report on semiautomatic segmentation using the Adobe Photoshop. *Journal of Digital Imaging* **18**, 333–343 (2005) [3](#)
24. Pielawski, N., Wählby, C.: Introducing Hann windows for reducing edge-effects in patch-based image segmentation. *PloS one* **15**(3), e0229839 (2020) [5](#)
25. Pollastri, F., Parreño, M., Maroñas, J., Bolelli, F., Paredes, R., Ramos, D., Grana, C.: A Deep Analysis on High Resolution Dermoscopic Image Classification. *IET Computer Vision* **15**(7), 514–526 (October 2021) [2](#)
26. Porrello, A., Vincenzi, S., Buzzega, P., Calderara, S., Conte, A., Ippoliti, C., Candeloro, L., Di Lorenzo, A., Dondona, A.C.: Spotting insects from satellites: modeling the presence of *Culicoides imicola* through Deep CNNs. In: *2019 15th International Conference on Signal-Image Technology & Internet-Based Systems (SITIS)*. pp. 159–166. *IEEE* (2019) [2](#)
27. Reyes-Herrera, P.H., Ficarra, E.: Computational Methods for CLIP-seq Data Processing. *Bioinformatics and Biology Insights* **8**, BBI-S16803 (2014) [2](#)
28. Ronneberger, O., Fischer, P., Brox, T.: U-Net: Convolutional Networks for Biomedical Image Segmentation. In: *Medical Image Computing and Computer-Assisted Intervention – MICCAI 2015*. pp. 234–241. *Springer* (2015) [3](#), [6](#), [9](#)
29. Sun, J., Darbehani, F., Zaidi, M., Wang, B.: SAUNet: Shape Attentive U-Net for Interpretable Medical Image Segmentation. In: *Medical Image Computing and Computer Assisted Intervention – MICCAI 2020*. pp. 797–806. *Springer* (2020) [6](#)
30. Weissheimer, A., De Menezes, L.M., Sameshima, G.T., Enciso, R., Pham, J., Grauer, D.: Imaging software accuracy for 3-dimensional analysis of the upper airway. *American Journal of Orthodontics and Dentofacial Orthopedics* **142**(6), 801–813 (2012) [3](#)
31. Zhou, Z., Rahman Siddiquee, M.M., Tajbakhsh, N., Liang, J.: Unet++: A nested u-net architecture for medical image segmentation. In: *Deep Learning in Medical Image Analysis and Multimodal Learning for Clinical Decision Support*. pp. 3–11. *Springer* (2018) [9](#)