

Knee Bone Segmentation on Three-Dimensional MRI

Rania Almajalid^{1,2}, Juan Shan^{*1}, Maolin Zhang³, Garrett Stonis⁴, and Ming Zhang^{*4,5}

¹ Department of Computer Science, Seidenberg School of CSIS, Pace University, New York City, NY, USA

² College of Computing and Informatics, Saudi Electronic University, Riyadh, Saudi Arabia

³ Hopewell Valley Central High School, Pennington, NJ, USA

⁴ Department of Computer Science & Networking, Wentworth Institute of Technology, Boston, MA, USA

⁵ Division of Rheumatology, Tufts Medical Center, Boston, MA, USA

ra56319p@pace.edu, jshan@pace.edu, maolinzhang@hvrds.org, stonisg@wit.edu, zhangml@wit.edu

Abstract—Three-dimensional (3D) images are widely used in the medical field (e.g., CT, MRI). In osteoarthritis research, 3D magnetic resonance imaging (MRI) provides a noninvasive way to study soft-tissue structures including hyaline cartilage, meniscus, muscle, bone marrow lesion, etc. The measurement of those structures can be greatly improved by accurately locating the bone structure. U-net is a convolutional neural network developed for biological image segmentation using limited training data. The original U-net takes a single 2D image as input and generates a binary 2D image as output. In this paper, we modified the U-net model to identify the bone structure on 3D knee MRI, which is a sequence of multiple 2D slices. Instead of taking a single image as input, the modified U-net takes multiple adjacent slices as input. The output is still a single binary image which is the segmentation result of the center slice in the input sequence. By using 99 knee MRI cases, where each knee case includes 160 2D slices, the proposed model was trained, validated, and tested. The dice coefficient, similarity, and area error metrics rate were tallied to assess the performance and the quality of the testing sets. Without any post-processing of the images, the model achieved promising segmentation performance with the Dice coefficient (DICE) 97.22% on the testing dataset. To achieve the best performance, diverse models were trained using different strategies including different numbers of input channels and different input image sizes. The experiment results indicate that the incorporation of neighboring slices generated better segmentation performance than using the single slice. We also found that a larger image size (uncompressed) corresponds to better performance. In summary, our best segmentation performance was achieved using five adjacent neighbor slices (two left neighbors + two right neighbors + the center slice) with the original image size of 352×352 pixels.

Keywords—knee osteoarthritis; 3D MRI images; deep learning; U-net; convolutional neural networks; automatic bone segmentation.

I. INTRODUCTION

Knee osteoarthritis (OA) is the most recurrent type of arthritis affecting the elderly. This form of arthritis is majorly responsible for handicaps and restrictions to carry out activities by elderly people. Research carried out in 2000 indicated that 13% of the population in the U.S. was 65 years old and above, and 50% of this population were affected by OA in at least one of their joints [1]. By 2030, it is approximated that 20% of the U.S. population, about 70 million individuals, will be aged 65 and will be predisposed to OA [1]. According to data analyzed in 2004, it was determined that the U.S. used approximately \$336 billion, equal to 3% of its gross domestic product (GDP),

to cater to patients of arthritis [2]. The knee osteoarthritis has been responsible for higher economic expense burden on society in terms of high costs associated with joint replacement, people leaving the workforce early, and higher incidences of employees being absent from work [3].

OA has no known cause or remedy that may slow the advancement of the ailment [4]. Magnetic resonance imaging (MRI) can provide a noninvasive way to study soft-tissue structures including hyaline cartilage, meniscus, muscle, bone marrow lesion, etc. However, it is very time-consuming to measure those structures using MR images because each MRI scan contains dozens of 2D images. For example, manual segmentation of cartilage for only one knee in three-dimensional (3D) knee MRI may take up to 6 hours. In addition, medical image readers require vast and comprehensive training to precisely measure cartilage [5]. The bone structure takes an important position in knee MRI and interacts with all other structures (cartilage, bone marrow lesion, meniscus, muscle). Bone segmentation from a knee MRI is essential to automatically measure other knee structures.

In the past decade, scientists have carried out investigations to get competent and stable methods to speed measuring knee structures such as bone, cartilage, and bone marrow lesion from the MR images. Some key strategies include restricting assessment to partial regions of those structures or segmenting alternate MR slices [6, 7, 8]. To automate the operations of segmentation for MR images, computer-aided algorithms based on B-splines and active contour-models have been exploited [6, 7, 9, 10, 11]. Nonetheless, these approaches were not reliable enough to be applied in clinical investigation, especially in exposing small structure changes [8]. Such a situation inspired the idea of identifying high accuracy structures first (i.e., bone) and then identifying other structures (cartilage, bone marrow lesion, etc.). Since the structure of cartilage is complicated, direct segmentation without bone identification is more challenging. Our bone identification study can serve as the first step for segmenting cartilage.

Deep learning (DL) methods, which are reputed for their potential to obtain high-level features, have proficiently taken care of critical problems in audio and vision fields [12, 13, 14]. DL methods have the ability to directly learn from the raw input, extract complex higher-level attributes layer by layer, and result in the exceptional achievement of classification and segmentation. Recently, enthusiasm over the application of DL approaches for medical image processing has increased [15]. U-

* Corresponding Authors: Juan Shan, jshan@pace.edu and Ming Zhang, zhangml@wit.edu

This research is supported by the National Science Foundation and Rheumatology Research Foundation.

net [16] is a DL method that has a distinctively U-shaped convolutional network architecture. Due to its precise and fast segmentation outcome, U-net, which was initially developed for neuron structure segmentation in microscopy images, won the ISBI challenge. U-net, being efficient at manipulating situations with minimal training datasets, is an exceptional match for our medical MR image segmentation.

In this work, we developed a fully automatic segmentation model for knee bone using a modified U-net structure. The modified U-net takes 3D information to improve its segmentation results. After a thorough investigation, we concluded that there is no prerequisite to tune any parameter for the dataset, as the method can be self-adjusted when used on a varying dataset.

The rest of the paper is organized as follows: the method is discussed in section II, including the modified structure of U-net. The experimental setup process including our dataset, the generation of the training and testing sets as well as the implementation are reviewed in section III. Section IV reports the experiment results while the conclusion is drawn in section V.

II. METHOD

A. Deep Convolutional Networks

Convolutional neural network (CNN) [17], is a deep learning architecture that directly derives features from image pixels as well as demands basic preprocessing. Recent exploration in computer vision as well as pattern recognition, has established that the CNNs have the proficiency to fix critical tasks including segmentation, object detection, and classification, with state-of-the-art performances [18, 19]. CNNs are comprised of subsampling and convolutional layers. They possess the potential to pinpoint patterns that cannot be identified by hand-crafted features. In the majority of instances, when CNNs are

supplied with adequate labeled data, they efficiently generate an outstanding hierarchical representation of the raw input images and produce an exemplary performance on computer vision tasks. Nonetheless, when CNNs are utilized to fix issues related to medical images, the meager quantity of images is a stumbling block for equipping a good model.

B. U-net

U-net [16] is a unique convolutional neural network architecture intended for exact and quick segmentation tasks. The U-net's architecture ordinarily is comprised of left and right paths which are the expansive and the contracting as shown in [16]. As illustrated in Fig. 1, the left contracting path follows the commonplace convolution network architecture. It consists of two convolution layers with a filter size of 3x3. Each layer is fortified by a 2x2 max pooling operation that has stride 2 which is utilized for down-sampling. A rectified linear unit (ReLU) follows every layer. On the other hand, the extensive path on the right contains three major components which are: up-sampling layer, feature map concatenation, and two 3x3 convolutional layers. Finally, a 1x1 convolution layer is used to map the 64-dimensional feature vectors into the pre-determined quantity of classes as a final layer in order to generate the segmented output.

Fig. 2 illustrates our inputs to the first layer; we modified U-net architecture to accept multiple channels as input (3D information) instead of one channel in the original U-net design. Since Knee MRI is a continuous scan (only 0.7mm between two slices), we expect that using the neighboring slices along with the processed slice can improve the segmentation result as more supporting information is provided. In our experiments, we performed the slice replication for slices near the two ends. To process the slice #2, we used the slices #1, #1, #2, #3, and #4 to form the five channels input, by duplicating the slice at the left side of the processed slice. Besides changing the number of the input channels, several other modifications have been done

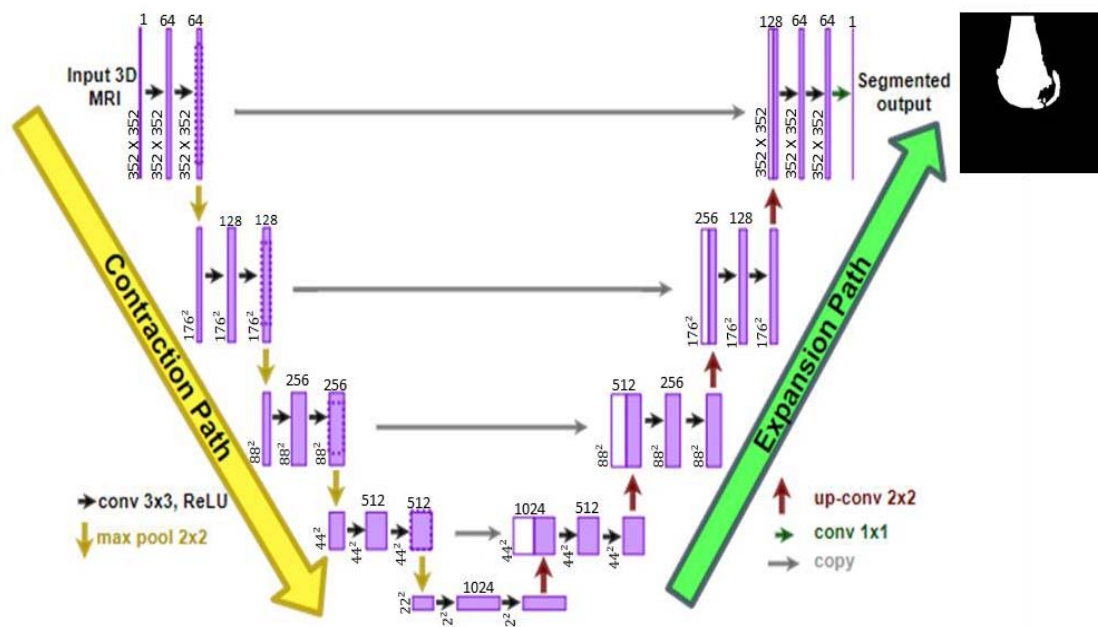


Fig. 1 Illustration of the modified U-net architecture.

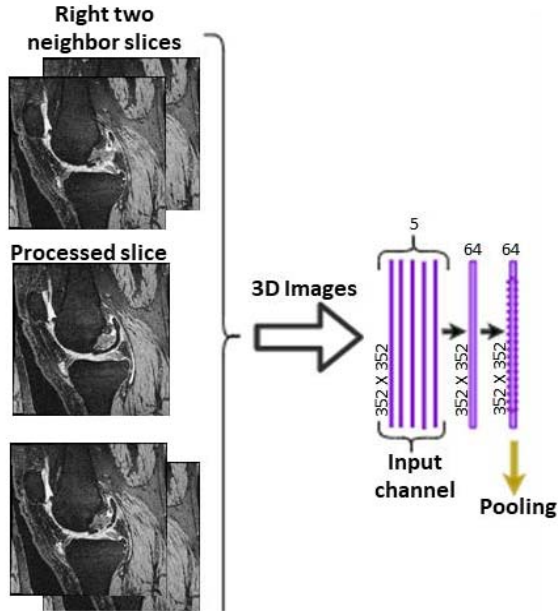


Fig. 2 Illustration of the inputs to the first layer.

on the original U-net. First, the original paper [16] did not do padding in any convolution layer so the pixels near the borders are lost after every convolution. Our model applied the padding to control the shrinkage of the image size after applying every convolution. Therefore, our output is the same size as the original input. Second, the original paper used the stochastic gradient descent optimizer, while we used the Adam optimizer which is a better optimization method adopted in many recent models [20]. Finally, the original paper used softmax with cross-entropy as the loss function. However, we used a soft DICE (minus DICE) as the loss function, which is a more direct loss measurement help to find the model with the best segmentation performance. The modified network requires 31,030,593 parameters to be equipped and consists of a total of 23 convolutional layers.

III. EXPERIMENTAL SETUP

A. Dataset

Our database is composed of 99 cases of 3D knee MRI sequences (15,840 DICOM images in total) that were obtained from the public OAI database and include all OA severity levels [21]. A single case consists of 160 2D slices with an original size of 384×384 pixels. We manually marked all the femur bones as ground truth. Since this study focuses on the segmentation of femur bone, all the slices that have a femur bone were selected. The selection has been done by finding the starting slice in which the femur bone started to show as well as the ending slice in which the bone disappeared, and all the images in between were used. Altogether, 11,701 DICOM images were selected from the 99 knee cases.

B. Generation of Training and Testing Sets

To prepare training and testing data for the U-net model, the 99 cases were divided randomly into three groups, train, test and validate sets, with each of the groups containing 70% (69 cases),

15% (15 cases), 15% (15 cases) of total knee cases respectively. The testing set has been held back until the end of the study. Be noted that the separation of the sets was done at the case level which means that all slice images from the same knee MRI scan should be placed in the same set. The training set for all the models contains 8200 slices (2D images) from 69 knees. The validation set contains 1708 slices from 15 knee cases while the testing set contains 1794 slices from another 15 knees.

Fig. 3(a) shows an example of the DICOM slices in our database. The femur bone was manually segmented for all slices, as appeared in Fig. 3(b). For every labeled DICOM slice, we generated a binary mask image by using a MATLAB script, as shown in Fig. 3(c). All original DICOM images were paired with its corresponding mask images. Furthermore, we cropped all our images 16 pixels from each side (right, left, top, and bottom) as shown in Fig. 3(d) to ensure that the bone starts from the very top edge of an image. As noticed in Fig. 3(b) that the software used for manual segmentation can't reach the very beginning of the images, we added this cropping operation to ensure the training data is accurate. After cropping, both the original image size and mask image size are changed from 384×384 pixels to 352×352 pixels.

C. Implementation

To implement our models, Keras [22] jointly with TensorFlow backend [23] in Python 3.7 was used. All experiments were conducted with a computer equipped with an NVIDIA GeForce GTX1080 Ti graphics processing unit (3584 GPU cores). All of our networks were trained using the Adam optimizer method that used the Dice coefficient (DICE) as a gauge for the accuracy of the segmentation process. On the other hand, the backpropagation through the CNN has been done using the soft DICE as a loss function. The size of the batches was set to 8. All the experiment models were set up for 300 epochs and using Keras's callback function called EarlyStopping, which terminates the training process of the model if the DICE has no improvement after a specified number of epochs, i.e. 20~30. This process saved a significant amount of time and it also avoided the model overfitting. Additionally, the rate of learning was initially set to 10^{-5} .

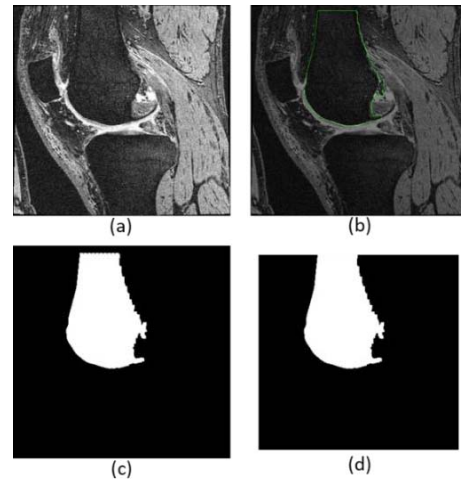


Fig. 3 (a) Raw image. (b) Manual segmentation. (c) Mask image generated from manual segmentation. (d) Mask image after cropping.

IV. EXPERIMENT RESULTS

A. Evaluation Metrics

The Dice coefficient (DICE) [24], also known as the overlap index is the most common metric used in authenticating medical image segmentation tasks. DICE is deduced by directly comparing the automatic segmentation and the manual segmentation results. It quantifies the spatial overlap rate between the two binary images. Its values vary between 0 and 1, with 1 representing a perfect match while 0 represents no overlap. Equation (1) depicts the DICE formula.

$$DICE = \frac{2 * |B_m \cap B_u|}{|B_m| + |B_u|} \quad (1)$$

In addition to DICE, we computed other area error metrics in order to evaluate our segmentation method from different perspectives. The similarity (SI), true positive (TP) ratio, false negative (FN) ratio, and false positive (FP) ratio, can be determined as below:

$$SI = \frac{|B_m \cap B_u|}{|B_m \cup B_u|} \quad (2)$$

$$TP \text{ ratio} = \frac{|B_m \cap B_u|}{|B_m|} \quad (3)$$

$$FN \text{ ratio} = 1 - TP \text{ ratio} \quad (4)$$

$$FP \text{ ratio} = \frac{|B_m \cup B_u - B_m|}{|B_m|} \quad (5)$$

The B_m refers to the pixel set of the manual segmentation of the bone outlined by an expert, while B_u refers to the bone region automatically generated by our trained model. Fig. 4 is an illustration of the area that corresponds to TP, FN, and FP. SI is a general proportion of the similarity between the contours of the automatic segmentation and the ground truth annotation.

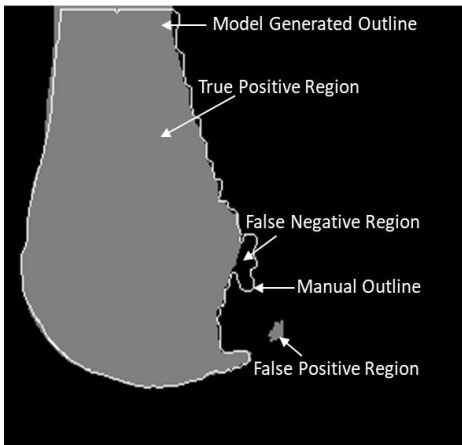


Fig. 4. Illustration of TP, FP and FN regions.

B. Experiments

To test the performance of the modified U-net structure on knee MRI, different models were trained with different strategies.

1) Different number of input channels

Since each 3D knee MRI contains a sequence of slices that represent the entire femur bone, we expect that using the neighboring slices along with the processed slice can improve the segmentation results. Different models were trained with different numbers of neighbors along with the processed slice. For the efficiency of the network training, all images and their corresponding masks were resized to 256 x 256 pixels. Table I summarized the performance of five models on the testing set. Model 1 is the original U-net which used one channel as input. Model 2 used three channels as input which were the processed slice along with its right and left slice neighbors. The output is the segmentation result of the processed (center) slice. The best performance was achieved by model 3 which had five channels as input, i.e., the processed slices along with its two right and two left slice neighbors. It reached 97.00% as an average DICE and 94.58% in term of an average similarity. Model 4 used seven channels as input and model 5 used nine channels. Overall, including more slice neighbors along with the processed slice outperformed the original U-net. Fig. 5 depicts two segmentation cases of the proposed model, which shows three examples for each case at different positions.

2) Different image size

The training of a deep learning model generally requires substantial computing power due to its extensive computation load to update millions of parameters in every single epoch. Therefore, the input image size greatly affects the training time. By compressing image size, one can greatly shorten the training time. In the previous experiment section, we shrunk the image size to save model training time. However, segmentation performance may be affected by image compression. In this section, we tested how image size affects U-net's segmentation performance.

Table II summarized the performance of our modified U-net on testing set with different input image sizes. The best model from last section with five channels was adopted here. We can clearly notice that in terms of an average DICE and the average similarity the results keep increasing while the image size increased to reach the original size of the images as 352 x 352 pixels. The best performance was achieved by the original image

TABLE I. THE PERFORMANCE OF TESTING SET OF FIVE U-NET MODELS WITH DIFFERENT INPUT CHANNELS NUMBER

Models*	TP (%)	FP (%)	FN (%)	SI (%)	DICE (%)
Model 1	95.58	02.81	04.42	93.73	96.35
Model 2	96.09	03.19	03.91	94.21	96.70
Model 3	96.14	01.89	03.86	94.58	97.00
Model 4	95.97	02.01	04.03	94.39	96.88
Model 5	95.72	01.81	04.28	94.38	96.89

* **Model 1** used exact one slice; **Model 2** used the slice and its 2 neighbors; **Model 3** used the slice and its 4 neighbors; **Model 4** used the slice and its 6 neighbors; **Model 5** used the slice and its 8 neighbors.

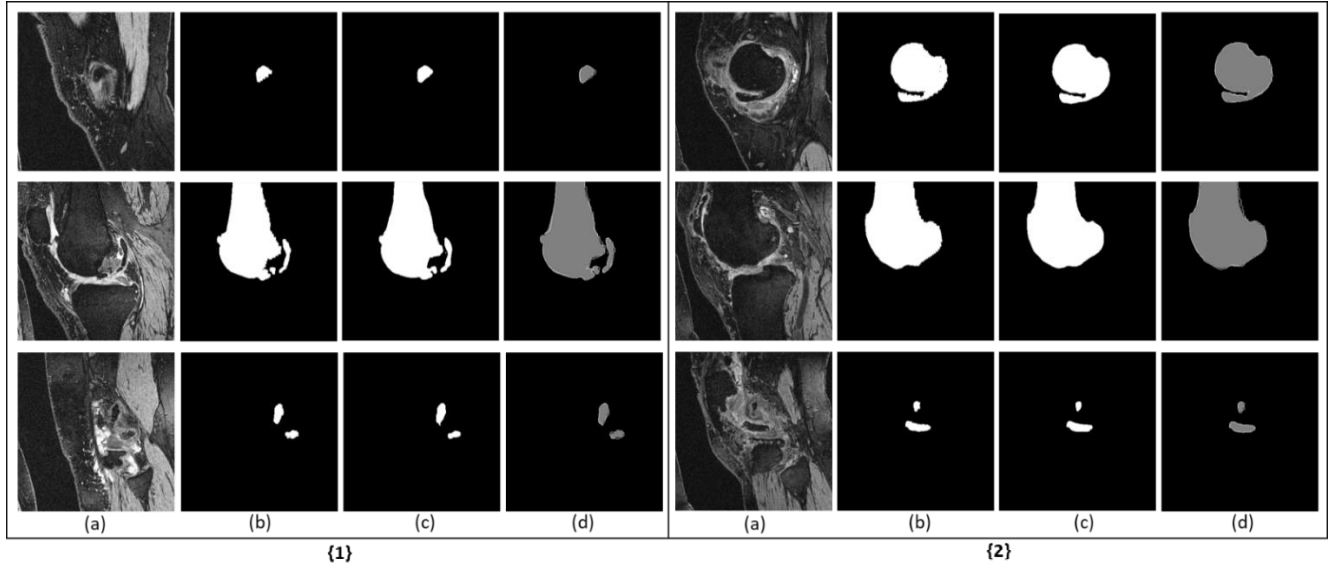


Fig. 5. (a) Processed slice. (b) Ground truth masks. (c) Output from proposed u-net. (d) Overlap between the automatic segmentation and the ground truth Note: Left three slices refer to the same case (knee) in different positions while the right three slices belong to another case.

TABLE II. THE PERFORMANCE OF TESTING SET OF PROPOSED U-NET MODEL WITH DIFFERENT IMAGE SIZES

Image Size	Time*	TP (%)	FP (%)	FN (%)	SI (%)	DICE (%)
128	1.40	93.41	01.33	06.59	92.51	95.83
192	2.54	94.93	01.42	05.07	93.82	96.56
256	4.26	96.14	01.89	03.86	94.58	97.00
320	6.23	96.38	01.92	03.62	94.73	97.03
352	7.67	97.13	02.43	02.87	94.98	97.22

*Time cost for training per epoch in minutes.

size and has 97.22% as an average DICE and 94.98% in terms of average similarity. This result coincides people's common sense that using original image size without compression can achieve the best accuracy, while we feel it is also meaningful to show the cost if one has to maintain smaller image size. In some situations of constrained computing resources or limited time on training, one could compress the image size to achieve a balance between accuracy and efficiency.

To further validate, we compared the performance of our proposed model along with original U-net longitudinally among different image sizes. As Table III showed, the proposed model outperformed the original U-net when using different image sizes in terms of femur-bone segmentation and achieved an average DICE = 97.22% and an average similarity = 94.98%, which is the best performance when using the original image size. Furthermore, in order to determine if there is a significant difference between the results of the original U-net and the proposed model, the student's t-test has been conducted in terms of DICE and similarity. The t-test results indicate that there is a significant difference at the $p = 0.05$ significance level for all models in terms of DICE and similarity. Table III illustrated the p -value for all the experiments. In addition, Fig. 6(a) provides a visual comparison between the original U-net and the proposed model using different image sizes in terms of similarity score

while Fig. 6(b) provides the same visual comparison in term of DICE score. The training time cost of each model is listed in Table III as well.

V. CONCLUSION

This study proposed a fully automatic segmentation pipeline based on U-net for bone segmentation on 3D knee magnetic resonance images. Without any image post-processing, the proposed model generated very remarkable segmentation results, i.e., DICE 97.22% and similarity 94.98% on the testing dataset. A comparison with the original U-net model was done and our proposed segmentation model outperformed it on the same dataset, with a statistically significant improvement. The results of the study established that providing neighboring slices in the 3D knee MRI sequence to U-net can improve the segmentation results of the femur bone in the knee joint.

One of our future work is to expand the experiment to the other bone structures in the knee joint, to include tibia and patella. This bone identification study can also serve as the critical step for segmenting other knee structures such as cartilage, bone marrow lesion, meniscus, and effusion. Since these structures are small and complicated, direct segmentation without bone identification is much more challenging.

TABLE III. THE COMPARISON OF THE MODIFIED U-NET MODEL WITH THE ORIGINAL U-NET MODEL USING DIFFERENT IMAGE SIZES

Image Size	Original U-net			The modified U-net			p -value (DICE)	p -value (SI)
	Time*	SI(%)	DICE(%)	Time*	SI(%)	DICE(%)		
128	1.20	91.29	95.09	1.40	92.51	95.83	0.0018	0.00004
192	2.20	92.77	95.98	2.54	93.82	96.56	0.0052	0.00013
256	3.43	93.73	96.35	4.26	94.58	97.00	0.0061	0.00338
320	5.40	93.81	96.38	6.23	94.73	97.03	0.0112	0.00241
352	6.60	94.40	96.73	7.67	94.98	97.22	0.0433	0.03689

*Time cost for training per epoch in minutes.

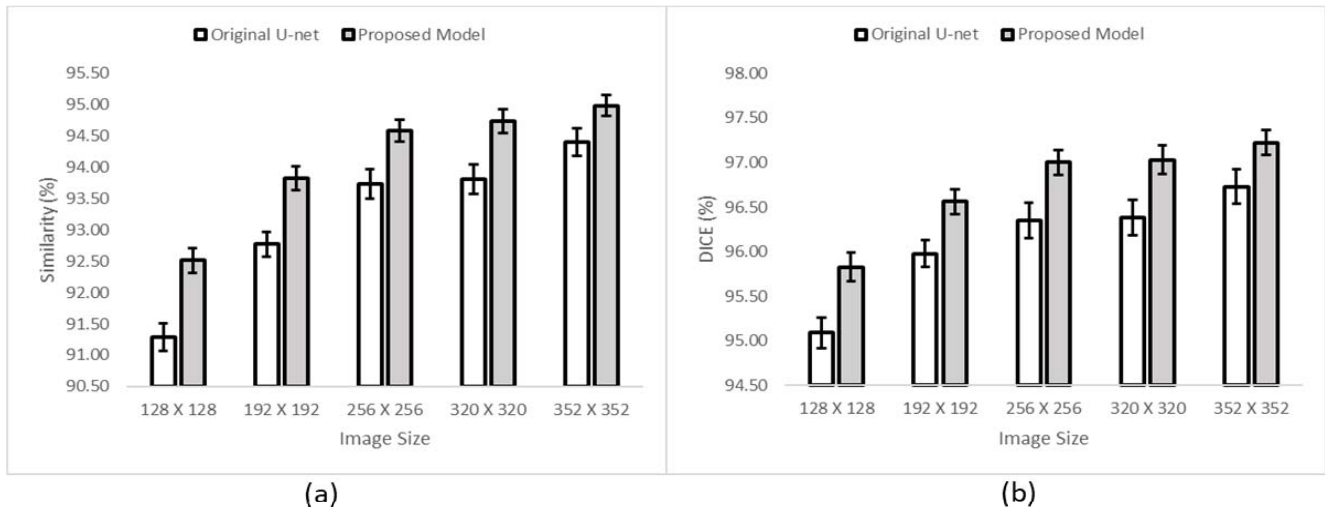


Fig. 6 The comparison of the performance of the proposed model with the original U-net among different image sizes in term of similarity (a) and DICE (b) results.

ACKNOWLEDGMENT

This research is supported by National Science Foundation awards (NSF-1723420, NSF-1723429) and Rheumatology Research Foundation award.

REFERENCES

- [1] National Institutes of Health, "Osteoarthritis Initiative Releases First Data," N. I. of Health et al., "Osteoarthritis initiative releases first data," NewsReleases, US Department of Health & Human Services, 2006.
- [2] U. S. Bone and J. Decade, "The burden of musculoskeletal diseases in the united states," Rosemont, IL: American Academy of Orthopaedic Surgeons, 2008.
- [3] A. A. Guccione, D. T. Felson, J. J. Anderson, J. M. Anthony, Y. Zhang, P. Wilson, M. Kelly-Hayes, P. A. Wolf, B. E. Kreger, and W. B. Kannel, "The effects of specific medical conditions on the functional limitation of elders in the framingham study," American journal of public health, vol. 84, no. 3, pp. 351–358, 1994.
- [4] Bhatia, D., Bejarano, T., and Novo, M.: 'Current interventions in the management of knee osteoarthritis', Journal of pharmacy & bioallied sciences, 2013, 5, (1), pp. 30–38
- [5] J. Jaremko, R. Cheng, R. Lambert, A. Habib, and J. Ronsky, "Reliability of an efficient mri-based method for estimation of knee cartilage volume using surface registration," Osteoarthritis and cartilage, vol. 14, no. 9, pp. 914–922, 2006.
- [6] Y. Yin, X. Zhang, R. Williams, X. Wu, D. D. Anderson, and M. Sonka, "Logismos layered optimal graph image segmentation of multiple objects on surfaces: cartilage segmentation in the knee joint," IEEE transactions on medical imaging, vol. 29, no. 12, pp. 2023–2037, 2010.
- [7] J. Fripp, S. Crozier, S. K. Warfield, and S. Ourselin, "Automatic segmentation and quantitative analysis of the articular cartilages from magnetic resonance images of the knee," IEEE transactions on medical imaging, vol. 29, no. 1, pp. 55–64, 2010.
- [8] F. Eckstein and W. Wirth, "Quantitative cartilage imaging in knee osteoarthritis," Arthritis, vol. 2011, 2010.
- [9] H. Z. Tameem and U. S. Sinha, "Automated image processing and analysis of cartilage mri: enabling technology for data mining applied to osteoarthritis," vol. 953, no. 1, pp. 262–276, 2007.
- [10] P. M. Cashman, R. I. Kitney, M. A. Gariba, and M. E. Carter, "Auto-mated techniques for visualization and mapping of articular cartilage in mr images of the osteoarthritic knee: a base technique for the assessment of microdamage and submicro damage," IEEE transactions on nanobioscience, vol. 99, no. 1, pp. 42–51, 2002.
- [11] G. Vincent, C. Wolstenholme, I. Scott, and M. Bowes, "Fully automatic segmentation of the knee joint using active appearance models," Medical Image Analysis for the Clinic: A Grand Challenge, vol. 1, p. 224, 2010.
- [12] H. Lee, P. Pham, Y. Largman, and A. Y. Ng, "Unsupervised feature learning for audio classification using convolutional deep belief networks," in Advances in neural information processing systems, pp. 1096–1104, 2009.
- [13] H. Lee, R. Grosse, R. Ranganath, and A. Y. Ng, "Convolutional deep belief networks for scalable unsupervised learning of hierarchical representations," in Proceedings of the 26th annual international conference on machine learning, pp. 609–616, ACM, 2009.
- [14] Q. V. Le, W. Y. Zou, S. Y. Yeung, and A. Y. Ng, "Learning hierarchical invariant spatio-temporal features for action recognition with independent subspace analysis," in Computer Vision and Pattern Recognition (CVPR), 2011 IEEE Conference on, pp. 3361–3368, IEEE, 2011.
- [15] A. A. Cruz-Roa, J. E. A. Ovalle, A. Madabhushi, and F. A. G. Osorio, "A deep learning architecture for image representation, visual interpretability and automated basal-cell carcinoma cancer detection," in International Conference on Medical Image Computing and Computer-Assisted Intervention, pp. 403–410, Springer, 2013.
- [16] O. Ronneberger, P. Fischer, and T. Brox, "U-net: Convolutional networks for biomedical image segmentation," in International Conference on Medical image computing and computer-assisted intervention, pp. 234–241, Springer, 2015.
- [17] Y. LeCun, L. Bottou, Y. Bengio, and P. Haffner, "Gradient-based learning applied to document recognition," Proceedings of the IEEE, vol. 86, no. 11, pp. 2278–2324, 1998.
- [18] R. Girshick, J. Donahue, T. Darrell, and J. Malik, "Rich feature hierarchies for accurate object detection and semantic segmentation," in Proceedings of the IEEE conference on computer vision and pattern recognition, pp. 580–587, 2014.
- [19] A. Krizhevsky, I. Sutskever, and G. E. Hinton, "Imagenet classification with deep convolutional neural networks," in Advances in neural information processing systems, pp. 1097–1105, 2012.
- [20] D. P. Kingma and J. Ba, "Adam: A method for stochastic optimization," arXiv preprint arXiv:1412.6980, 2014.
- [21] Imorphics. [Online]. Available: <http://imorphics.com/>. [Accessed: 28-Aug-2019].
- [22] F. Chollet, "Keras," <https://github.com/fchollet/keras>, 2015.
- [23] M. Abadi, A. Agarwal, P. Barham, E. Brevdo, Z. Chen, C. Citro, G. Corrado, A. Davis, J. Dean, M. Devin, et al., "Tensorflow: Large-scale machine learning on heterogeneous distributed systems," 2016.
- [24] L. R. Dice, "Measures of the amount of ecologic association between species," Ecology, vol. 26, no. 3, pp. 297–302, 1945.