

A Multi-Scale Fully Convolutional Networks Model for Brain MRI Segmentation

Zhihui Cao, Yuhang Qin, Yunjie Chen*

¹*School of math and statistics, Nanjing University of Information Science & Technology, Nanjing 210044, China*
(Received October 01 2017, accepted January 15 2018)

Abstract. Accurate segmentation for brain magnetic resonance (MR) images is of great significance to quantitative analysis of brain image. However, traditional segmentation methods suffer from the problems existing in brain images such as noise, weak edges and intensity inhomogeneity (also named as bias field). Convolutional neural networks based methods have been used to segment images; however, it is still hard to find accurate results for brain MR images. In order to obtain accurate segmentation results, a multi-scale fully convolution networks model (MSFCN) is proposed in this paper. First, we use padding convolutions in conv-layer to preserve the resolution of feature maps. So we can obtain segmentation results with the same resolution as inputs. Then, different sized filters are utilized in the same conv-layer, after that, the outputs of these filters are concatenated together and fed to the next layer, which makes the model learn features from different scales. Both experimental results and statistic results show that the proposed model can obtain more accurate results.

Keywords: Convolutional neural networks; Fully convolutional networks; magnetic resonance image; multi-scale.

1. Introduction

Brain disease is one of the main diseases that threaten human health, so it is of important sense to use brain imaging to help us diagnose the brain disease [1-3]. Compared with other medical images, brain MR images are easier to be used for diagnosis of brain disease for their high contrast among different soft tissues and high spatial resolution [4]. Segmenting brain tissues accurately, including white matter (WM), gray matter (GM) and cerebrospinal fluid (CSF), plays an important role in both clinical practice and medical study. However, some imaging artifacts such as noise, intensity inhomogeneity and weak edges, which drives scholars to find and propose more robust and more accurate approaches, hinder most segmentation methods.

Quite a lot of clustering algorithms have been proposed for brain MR image segmentation. Fuzzy C-means (FCM) algorithm, first introduced by Dunn [5], is one of the most widely used ones. FCM assumes that a pixel of an image belongs to different classes at different degree, corresponding to brain MR images' fuzziness. The FCM fails to segment images with noise, low contrast and bias field by only using intensity information. In order to improve the robustness of FCM, many scholars proposed modified models based on FCM by adding spatial information into it and obtained a certain improvement [6, 7]. However, they failed to solve the problem of low contrast.

In the last several years, deep learning (DL) [8], especially convolutional neural networks (CNN), has outperformed the state of the art in computer vision tasks. Since the breakthrough by Krizhevsky et al. [9], even larger and deeper networks have been trained [10, 11].

The traditional use of CNN is on classification tasks, where the output we want is just a class label for the input image. However, in image segmentation tasks, the desired output should be an image with each pixel labeled. Hence, Cirosan et al. [12] trained a network (DNN) to predict a class label of a pixel by labeling a patch in a square window centered on the pixel itself. They succeeded to apply CNN to image segmentation and won the EM segmentation challenge at ISBR 2012. Nevertheless, there are some shortcomings in this model. First, it requires a great lot of calculation and room, for example, if we want to segment an image with a size of 512×512 , we have to classify 262144 patches with the network. Secondly, there is a lot of redundancy because of the overlapping patches of adjacent pixels. Thirdly, there is a tension between local information and global information. Small patches guarantee the localization accuracy but use little context, while large patches may reduce the localization accuracy. These cause the network inefficient.

Long et al. [19] first trained an end-to-end learning network for semantic segmentation called fully convolutional networks (FCN), which outputs dense prediction from arbitrary-sized inputs. Moreover, both

training and predicting are performed whole-image-at-a-time. In-network convolutional layers extract features and upsampling layers enable pixelwise prediction. Compared with patch-wise training networks [12], the FCN model is more uncomplicated and works more efficient with no pre- or post-processing.

However, there are millions of parameters to be trained in the FCN model, which requires a large amount of training data and a lot of training time, but it is unrealistic in biomedical tasks. Hence, Ronneberger et al. [20] modified and extended the FCN model such that it works with few training images and obtains more precise results, and this model is called U-Net for its U-shaped architecture. U-Net upsamples at stride 2 and combines shallow layers with upsampled outputs, thus it learns better and the segmentation results are more realistic with sharper edges. But U-Net underperforms when it comes to segmenting details, such as CSF in brain MR images.

Based on the analysis above, we can find gray value based methods are sensitive to outliers. Although some modified models reduce the influence of noise and outliers to some extent, most improvements are at the price of increasing parameters and complexity of the model. FCN and U-Net are efficient and find better results, while details and narrow bands in images are lost. To address these drawbacks, we propose a multi-scale FCN model (MSFCN) in this paper. Different-scaled filters are added into the model to obtain more accurate results, where large-scaled filters see more context and small-scaled ones keep details and narrow bands. We have compared proposed model with other state-of-the-art segmentation models to show that our model can obtain more precise results.

2. Backgrounds

2.1 Deep convolutional networks (DNN) in segmentation

Considering the excellent performance of deep learning in classification task, many scholars try to take advantage of it for image segmentation. Patchwise training was used in many approaches, in which each pixel is labeled with the class of its enclosing region. For example, Ciresan et al. [12-18] succeeded to apply CNN to segmentation task, but the poor efficiency made it impossible for their model to be used in clinical medicine.

By contrast, FCN [19] is more efficient. Long et al. reinterpreted classification networks as fully convolutional and add upsampling layers, which decreased parameters and complexity a lot and enabled pixelwise prediction. However, there are some drawbacks in the FCN model. First, the FCN model takes some typical CNN models such as AlexNet, VGGNet, etc as its contracting part, which requires a large amount of training data. Secondly, the model has to be trained three times (FCN-32s, FCN-16s, FCN-8s). Thirdly, the result of FCN is too smooth and the details are lost.

Based on FCN, Ronneberger et al. [20] built a more elegant architecture called U-Net. This network uses successive convolutional layers followed by a maxpooling layer in contracting part, and in expansive part, the process is inverted. Further, feature maps with the same resolution from both parts are concatenated together followed by successive convolutional layers and two 1×1 filters are used in the final to obtain the segmentation results. But due to the unpadded convolution, the resolution of results is lower than that of inputs. Moreover, so many times of convolutions make the details lost in high layers, leading to erroneous results at these pixels.

3. Proposed Model

We replace some 3×3 kernels with 1×1 and 5×5 ones based on U-Net without changing the depth of the network, which enable the network to extract features from different scales at the same time. In the following, an overview of our model is given.

3.1 Motivation

Szegedy et al. [10] trained GoogLeNet model in 2014 and won the classification task of Imagenet Large Scale Visual Recognition Challenge 2014(ILSVRC2014). They proposed an inception module, where 1×1 , 3×3 , 5×5 filters and max-pooling are used at the same time. This work improved their final accuracy and reduced the number of parameters quite a lot.

Considering the improvement by the inception module, we proposed MSFCN model for segmenting brain MR images. Different from classification task, a class label should be assigned to each pixel in segmentation task, which means not only the global information and large regions matter, but small regions, details and the pixel itself are also essential. Thus, we increase the proportions of 1×1 and 3×3 filters and