

上海交通大学

项目报告

实习单位: University of Illinois Urbana-Champaign

实习时间: July 23th 至 Sept 7th

学院(系): Biomedical Engineering

专 业: Biomedical Engineering

学生姓名: Chaowei Wu 学号: 515021910488

2018 年 08 月 30 日

New Strategy for Super-resolution in Magnetic Resonance Imaging

Chaowei Wu 515021910488

Abstract

Background

The goal of our super resolution project is to generate high resolution images, which recover subjects' unique features and remove artifacts or ringing, from down-sampled low-resolution MRI images.

Method

As the very first step, we utilize convolution neural network to provide us a very good initial estimation of the high-resolution data. And as the next step, we aim to extract and recover novel features which existed in practical patient images, based on k-space residual and sparsity constraint. As the ultimate step, we utilize optimization model based on data consistency to eliminate artifacts and give the final prediction.

Results

Based on our method, we convert every single MPRAGE image to 4 times high resolution image with well recovered novel feature and little artifact.

Background

1.1 Magnetic Resonance Imaging (MRI)

Magnetic Resonance Imaging has been more and more widely used today in disease diagnosis for diseases like brain diseases, spinal disorder, cardiac function and so on. Although MRI requires longer requisition time than Computed Tomography (CT), it doesn't require ionizing radiation and can be performed in any orientation. MRI image

has excellent 3-D capabilities, brilliant tissue contrast and high spatial resolution. [1]

1.2 Super-resolution Technique

However, there always exists trade-off problem between the resolution and speed in MRI. Limited by acquisition time and motion of subject, well-sampled image sometimes is hard to get. Instead, faster imaging techniques, such as Parallel Imaging [2], and post-acquisition processing techniques, such as Super-resolution [1] and Motion Correction [3, 4], are getting concerned.

Super-resolution technique is to recover high-resolution images from down-sampled images. It offers higher resolution, ensuring additional significant information and increase the diagnosis possibilities.

Identified Problem

Super-resolution is an ill-posed problem and it's difficult to achieve. (Fig.1). It aims to recover k space with $k_{max} = k_1$ in the condition that only scan for $k_{max} = k_2 (k_2 < k_1)$. To fill in all the fields which are zero-padded, additional prior information is needed.

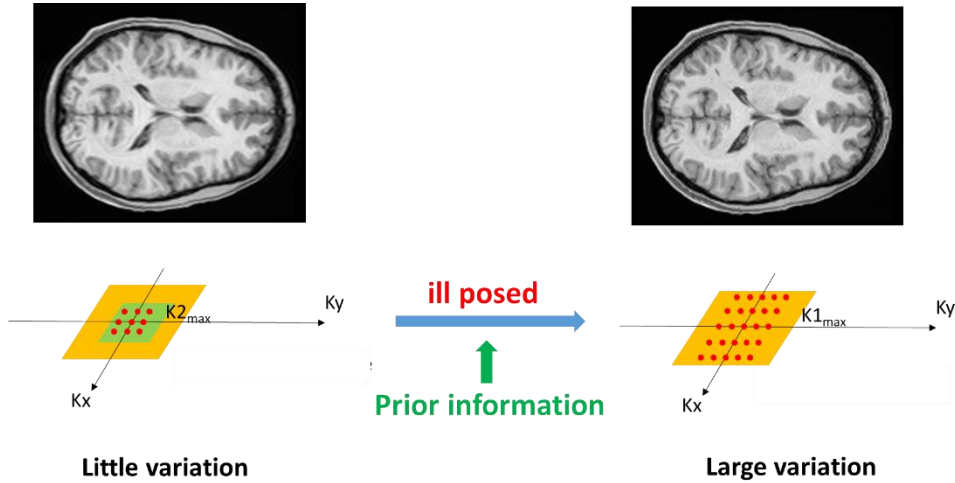


Fig.1 Ill-posed Problem in Super-resolution technique

To solve this problem there are two main kinds of methods. [5]

2.1 Regularized Formulation

One is regularized formulation to reconstruct from sparse or noisy measurements to a high-quality image.

$$R_{reg}\{y\} = \operatorname{argmin}_{x \in \mathcal{X}} f(H\{x\}, y) + g(x) \quad (1)$$

In which $H\{x\}$ is degradation model, f is cost function and $g(x)$ is regularization term that promotes solutions that match our prior knowledge of x and makes the problem well posed simultaneously. Though, for the sake of its simple form, and difficulty to find suitable degradation model, prior information it can absorb is limited.

2.2 Learning-based Method

The other solution, where a training set of ground-truth images and their corresponding measurements is known. A parametric reconstruction algorithm, R_{learn} , is then learned by solving

$$R_{learn}\{y\} = \operatorname{argmin}_{R_\theta, \theta \in \Theta} \sum_{n=1}^N f(x_n, R_\theta\{y_n\}) + g(\theta) \quad (2)$$

Θ is the set of all possible parameters, f is a measure of error, and $g(\theta)$ is a regularizer on the parameters to avoid overfitting. Once the learning step is complete, R_{learn} can then be used to reconstruct a new image from its measurements. In this way, a deep learning based method has been built. The neural network can be deep and complex, and it can absorb more prior information.

2.3 Problems in current learning-based methods

In the last few years, plenty of researchers have been working on Super-resolution.[6, 7] Nevertheless, they only focus on the satisfying visual resolution generously, but ignore artifacts and poorly recovered novel feature. (Fig.2)



Fig.2 Distorted pattern (Red) and novels (Blue)

This issue can be acceptable in natural pictures but will become fatal in medical images.

Medical images, in one sense, are created to find the novel features and help diagnosis. If medical image loses novel feature and contains such kind of artifacts, it will influence doctors' judgment about patients. So that's the problem we aim to solve.

Potential Solutions

We can briefly summary our methods as follow, or as shown in Fig.3.

- As the very first step, we utilize convolution neural network to provide us a very good initial estimation of the high-resolution data.
- And as the next step, we aim to extract and recover novel features which existed in practical patient images, based on k-space residual and sparsity constraint.
- As the ultimate step, we utilize optimization model based on data consistency to eliminate artifacts and give the final prediction.

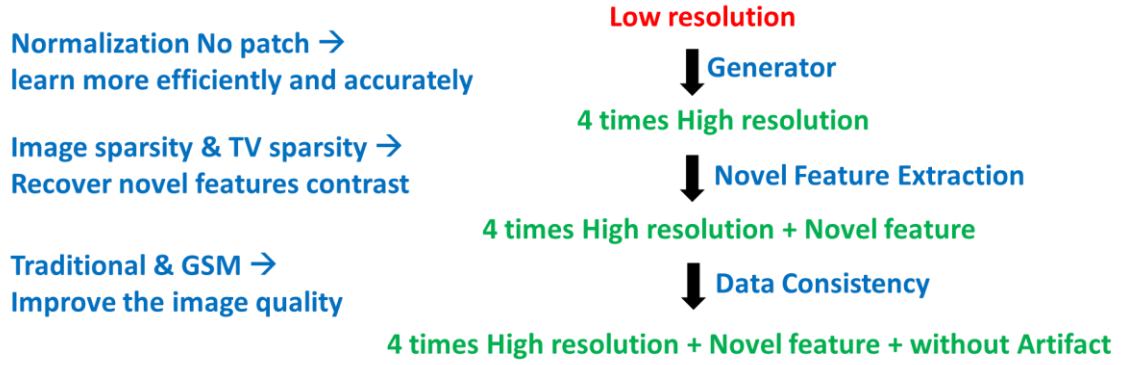


Fig.3 Preliminary Pipeline and Methods

3.1 Step#1 Deep Learning

As the very first step, we utilize convolution neural network to provide us a very good initial estimation of the high-resolution data.

3.1.1 Analysis of machine learning

Machine learning is to learn $P(\text{output} | \text{input})$, and give the most probable output when certain data inputs. We can express our image as a function of θ_t (tissue property), θ_p (parameters related to imaging) and θ_g (geometry of the brain). (Fig.4) In this case, machine learning is to learn the joint distribution of pixel value and $\theta_t, \theta_p, \theta_g$. [8]

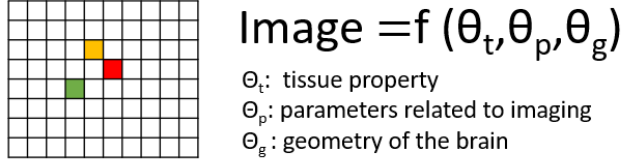


Fig.4 Image Model

By default, researchers train their model with same modality so that θ_p is constant.

In traditional way, researchers patch their images to throw them into network, then get output of each patch and concatenate them into a whole image. They also don't implement image normalization to reduce variation of θ_g to least. [6, 7] In this way, what the Convolution Neural Network (CNN) learns is $P(\theta_t - \text{high}, \theta_g - \text{high} | \theta_t - \text{low}, \theta_g - \text{low})$.

3.1.2 Design of Data Preprocessing

To reduce the difficulty of network training, we decided to leverage the advantage of image normalization which based on DARTEL algorithm.[9] In this way, we minimize the influence of θ_g (geometry variation). What's more, we don't patch data, which means whole images are thrown into neural network. By doing so, we aim to help CNN learn $P(\theta_t - \text{high} | \theta_t - \text{low})$ more accurately and efficiently.

3.1.3 Design of Neural Network Structure

We utilize U-net Generative Adversarial Network (GAN) as our CNN architecture, and Res-net GAN as a contrast. U-net GAN architecture is also utilized in [10] for super-resolution. These network architectures are shown below. (Fig.5,6,7)

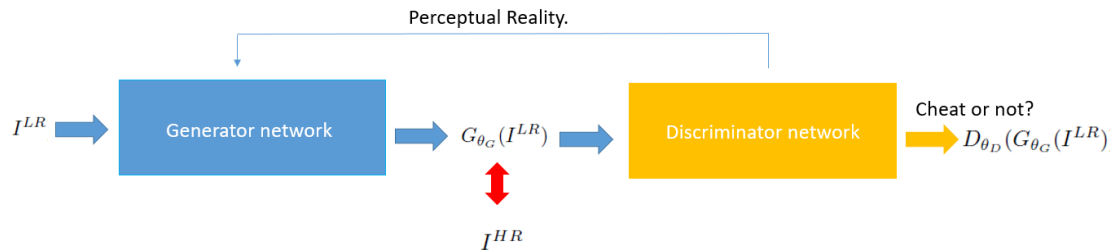


Fig.5 Generative Adversarial Network (GAN) architecture

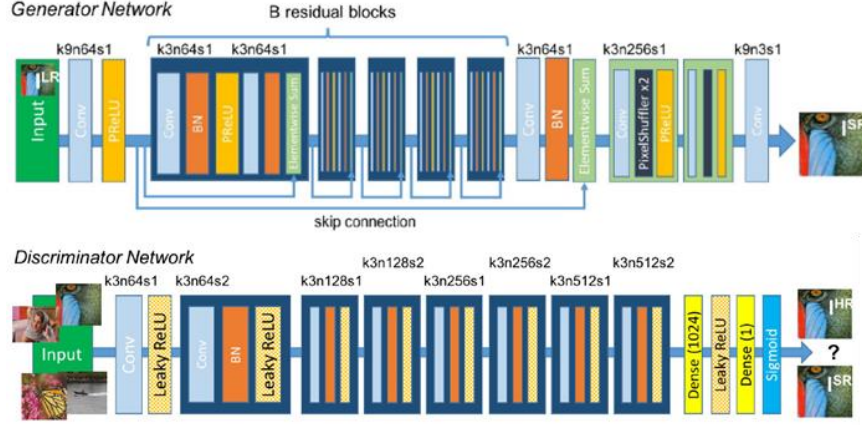


Fig.6 Resnet-GAN architecture

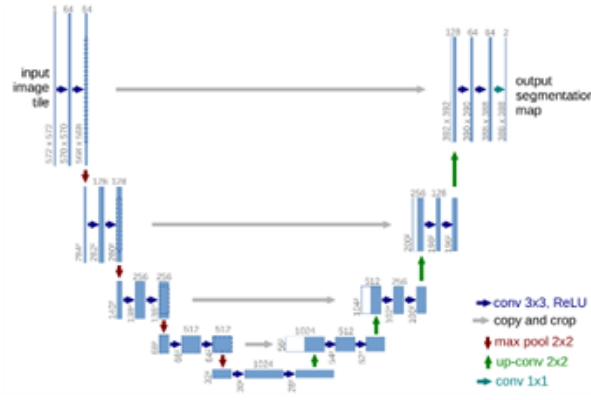


Fig.7 U-net architecture

3.2 Step#2 Novel Feature Extraction

The CNN output has much higher resolution than input, but the novel feature is still obscure. As the next step, we aim to extract and recover novel features which existed in practical patient images, based on k-space residual and sparsity constraint.

3.2.1 Assumption

We build our model based on two assumptions:

1. Difference in inner k space contains novel feature
2. Feature is sparse in some domain (Image domain, Total Variation (TV), etc)

3.2.2 Model

Our model names *Sparsity Based Novel Feature Extraction model*. This optimization model can be expressed as below:

$$\operatorname{argmin}_{\rho} ||d' - \Omega F \rho||_2^2 + \lambda ||W \rho||_1 \quad (3)$$

$$d' = d - d_{ref}$$

In which d is true inner k space data, d_{ref} is data consistency result k space data, W represents sparsity transform, Ω represents truncation operator, F is Fourier Transform Operator, λ is regularization constant, and ρ is to be solved.

In brief, we aim to extract features from difference of true inner k space and CNN prediction inner k space, as is shown in Fig.8.

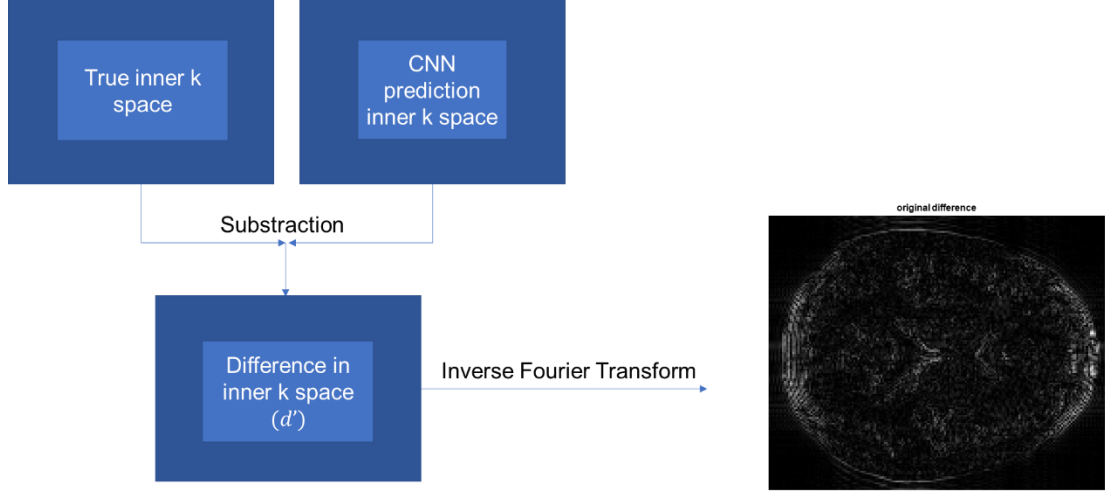


Fig.8 Procedure of getting d' in Step#2

3.2.3 Sparsity Transform

In our *Sparsity Based Novel Feature Extraction model*, we need specific sparsity transform W to transform feature into domains in which it's sparse. For now, we tried two, one is $W = I$ (identity matrix), and the other is $W = \nabla$ (Total Variation).

3.3 Step#3 Optimization based on Data Consistency

After novel feature extraction, the image still has some artifacts and ringing. As the ultimate step, we utilize optimization model based on data consistency to eliminate artifacts and give the final prediction. Here, we build two different models to solve it.

3.3.1 Traditional Data Consistency Model

We build this model based on Fourier Series. And we believe our ideal prediction should be a balance of input data and novel feature step recovered image.

$$\operatorname{argmin}_x ||y - \Omega Fx||_2^2 + \lambda ||P - x||_2^2 \quad (4)$$

In which y is true inner k space data, P is novel feature step recovered image,

Ω represents truncation operator, F is Fourier Transform Operator, λ is regularization constant, and x is to be solved.

3.3.2 Generalized Series Method (GSM) [11]

We believe artifacts can be eliminated by generalizing series basis to remove the variation in k space. The model can be expressed as:

$$\begin{aligned} \argmin_{c_l} ||d - \Omega \odot [\tilde{d}(k) * \sum_{l=-L}^L c_l \delta(k - l\Delta k)]||_2^2 + \lambda \sum_l ||c_l||_2^2 \quad (5) \\ \Rightarrow \argmin_{c_l} ||d - Ac||_2^2 + \lambda \sum_l ||c_l||_2^2 \end{aligned}$$

In which d is true inner k space, $\tilde{d}(k)$ is k space of novel feature recovered image, Ω is mask, L is kernel size, λ is penalty constant, and c_l is to be solved. Once c_l is computed, we can get our predicted image by $F^{-1}[\tilde{d}(k) * \sum_{l=-L}^L c_l \delta(k - l\Delta k)]$.

Results and Discussion

4.1 Step#1 Deep Learning

We obtain images with much higher resolution in this step with the help of image normalization and U-net GAN.

4.1.1 Image Normalization Effect

As shown in Fig.9, after normalization, both training loss and valid loss drops more rapidly. Moreover, valid loss of normalized data is much lower than that of non-normalized when it reaches convergence. Therefore, normalization does help learning converge faster and lower the loss.

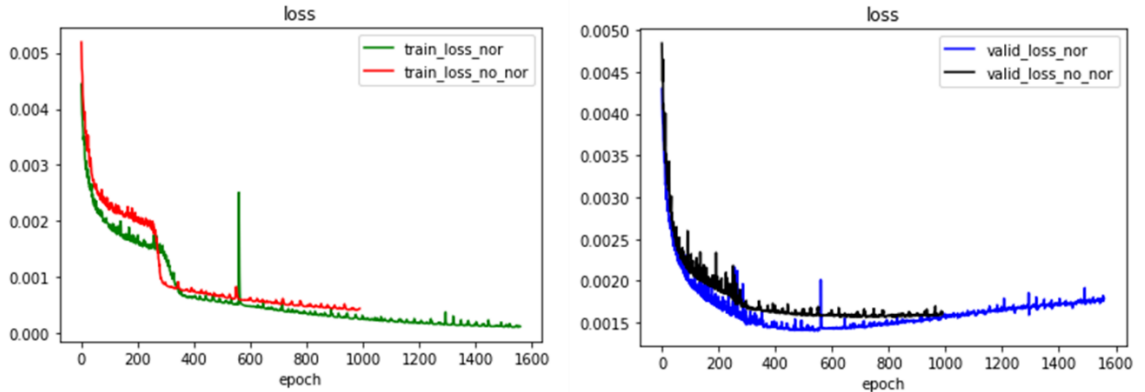


Fig.9 Training loss and Valid loss with normalized or non-normalized data

And as shown in Fig.10, prediction of CNN trained with normalized data has similar contrast with ground truth and doesn't generate new artifacts (Red frame).

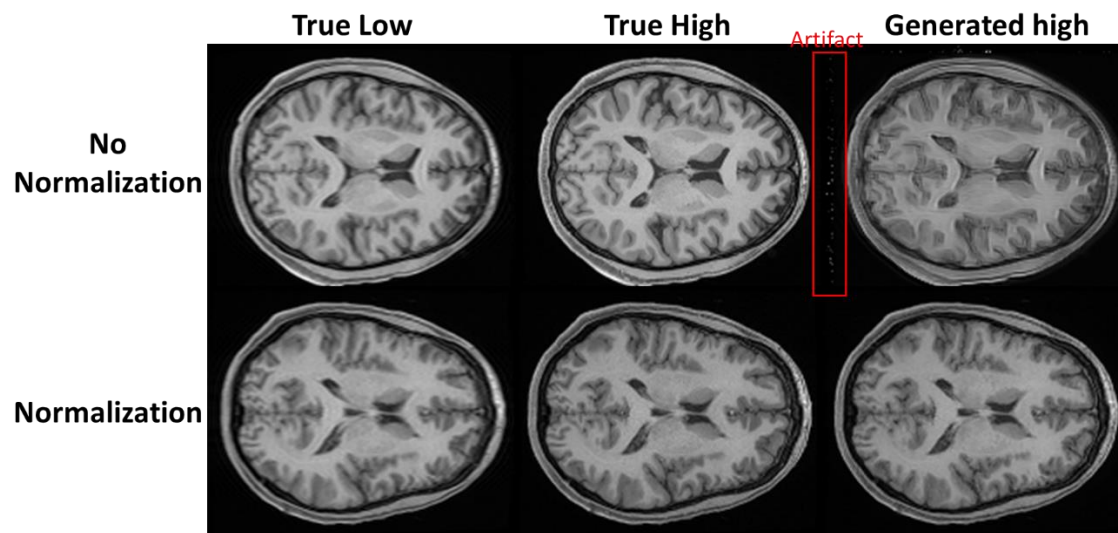


Fig.10 CNN Prediction – trained with normalized data or non-normalized data

4.1.2 Comparison between Resnet-GAN and U-net GAN

Basically, results of the two network architectures (Fig.11) are compared in three aspects: resolution, novel feature recovery and artifact.

1. **Resolution.** They are comparative in resolution.
2. **Novel Feature Recovery.** In Deep Learning step, we want CNN to absorb as much priori as possible, so in case we only train it with normal subjects, it should regard tumor as normal tissue, because in normal subjects there should be gray matter or white matter. Therefore, our ideal CNN output should have obscure or even no novel feature. In our results, tumor is rather light and clear in Resnet-GAN output, while it's obscure and dim in U-net GAN output, which proves that Resnet-GAN learns a local sharpening, while U-net GAN learns a distribution. It still needs further study.
3. **Artifact.** Res-net GAN generates more artifacts. In Fig.11, a drill is found in the red frame in Res-net GAN output, but it doesn't exist in both U-net GAN output

and ground truth. There's another example: a drill appears against expectation in Fig.12.

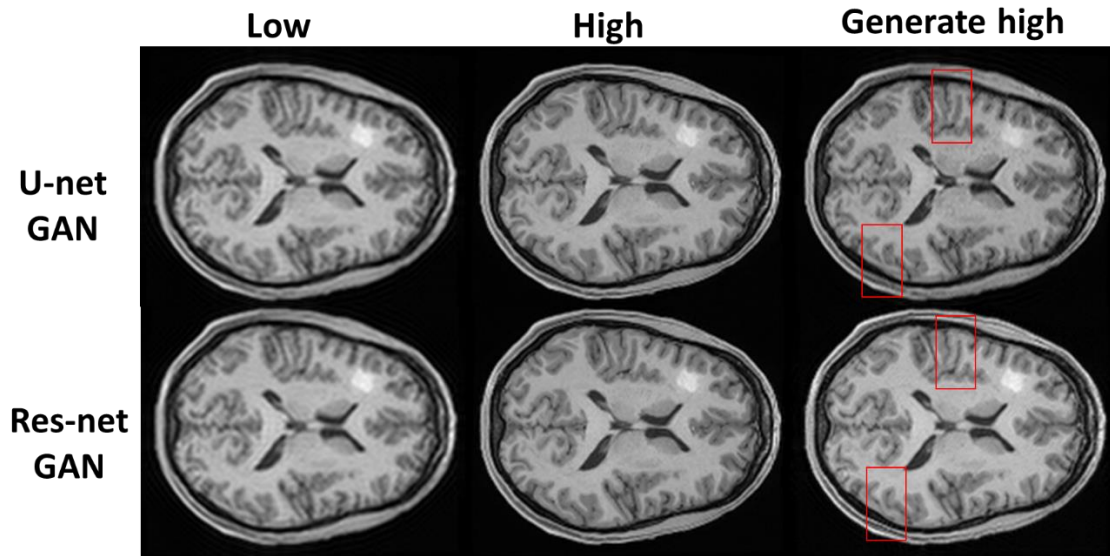


Fig.11 Comparison between U-net GAN result and Res-net GAN result

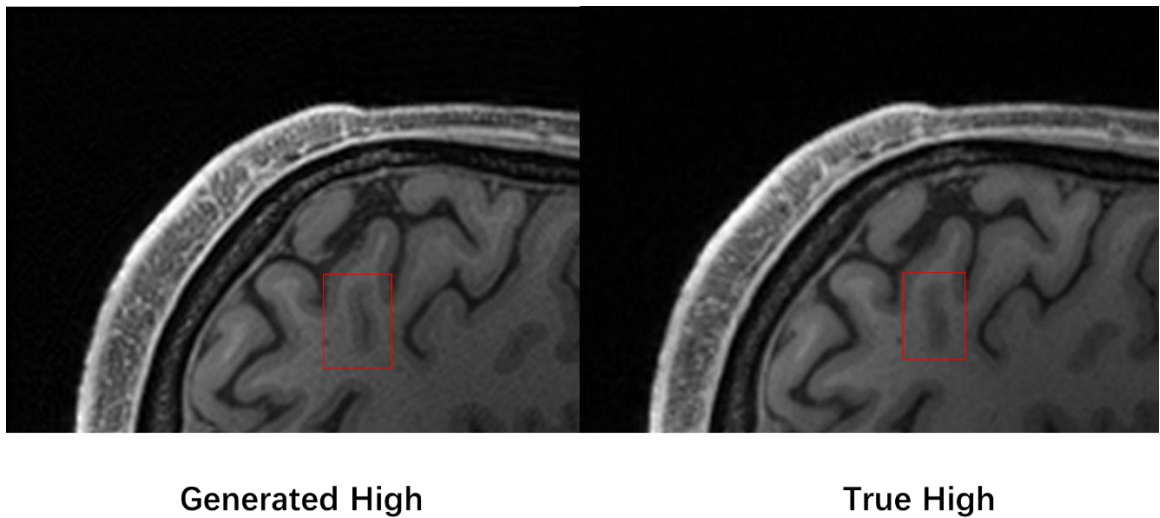


Fig.12 Artifacts generated by Resnet-GAN

4.2 Step#2 Novel Feature Extraction

After Deep Learning step, we get images with much higher resolution but poorly recovered novel feature (Fig.13). Then we extract novel feature with our model with sparsity constraint. We find that we could extract features and enhance its contrast, but the limitation is that it's still low-resolution.

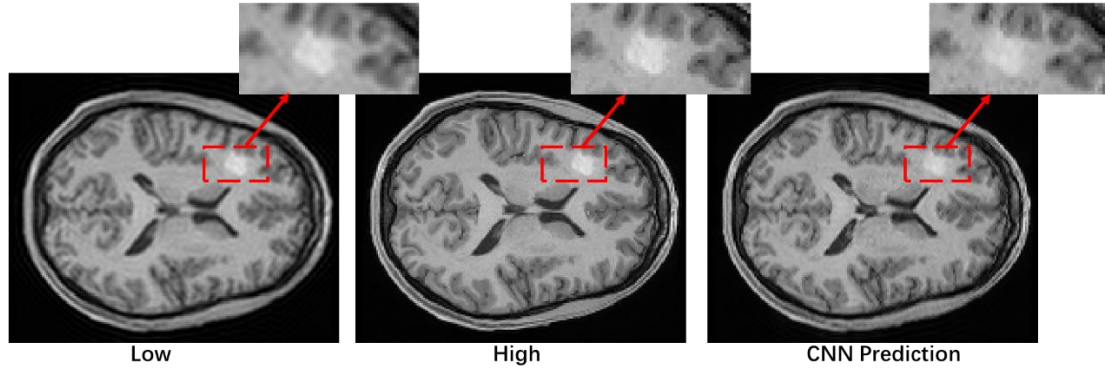


Fig.13 Poorly Recovered Feature in CNN

4.2.1 Image Domain Sparsity Result

As shown in Fig.14, we get difference in inner k space in Fig.8 procedure, and solve (3) model, and get our extracted novel feature. We can notice that pixel values in tumor are a little lighter. Then we get our feature-recovered image by procedure in Fig.14.

We can discover enhanced contrast but still low resolution in Fig.15. Also, because extracted novel feature contains ringing in some degree, we need to implement Step#3, optimization based on data consistency.

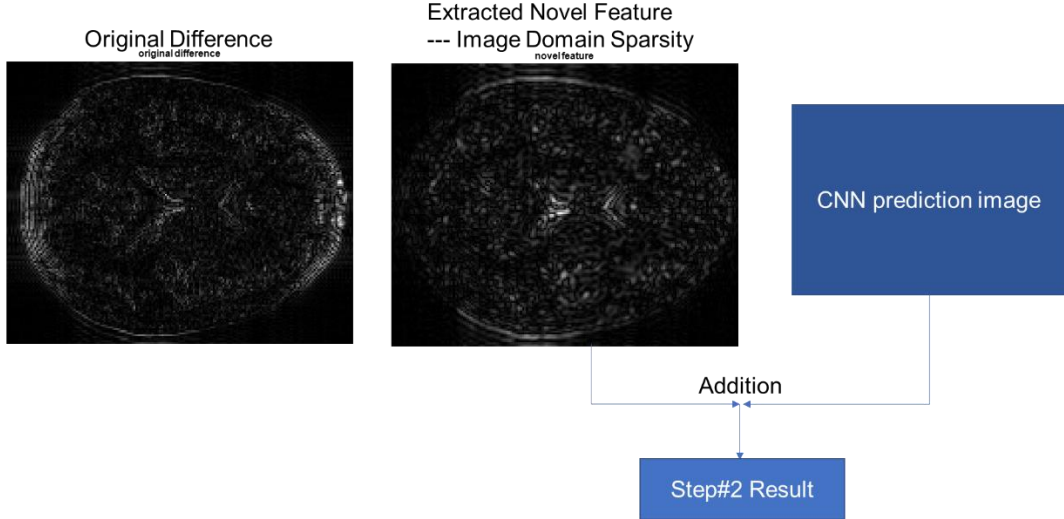


Fig.14 Extracted novel feature with image domain sparsity

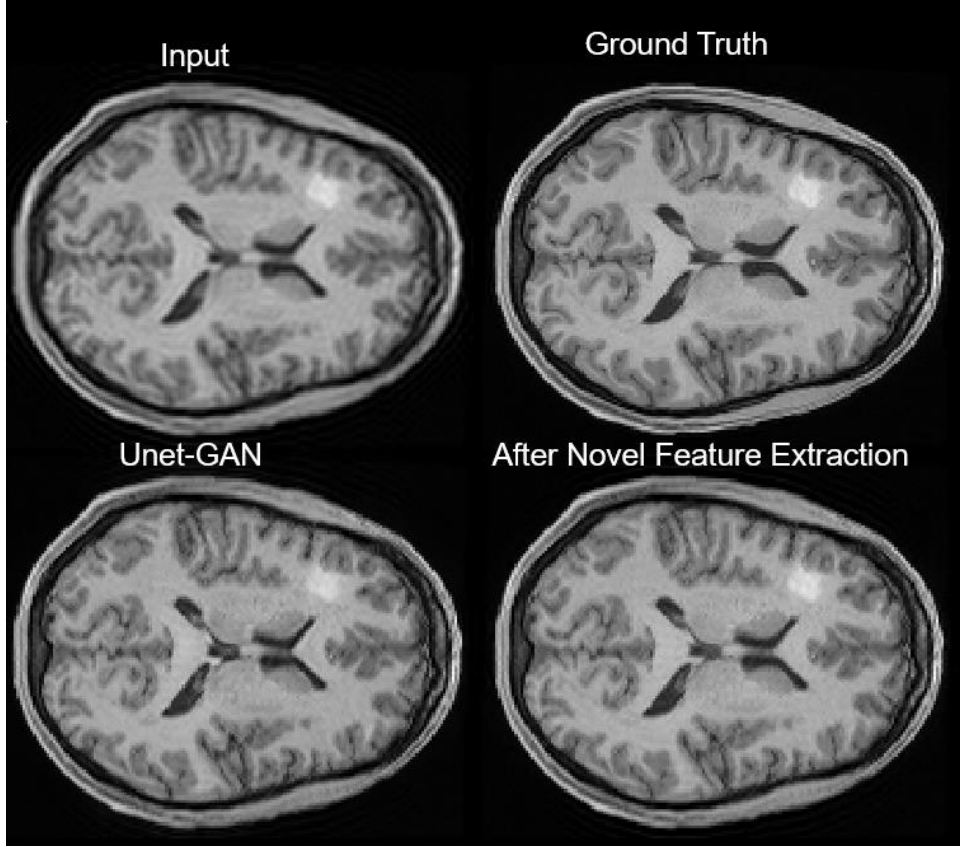


Fig.15 Novel Feature Extraction result with image domain sparsity constraint

4.2.2 Total Variation Sparsity Result

The Total Variation (TV) in one direction is defined as:

$$V(y) = \sum_n |y_{n+1} - y_n| \quad (6)$$

We utilize TV in two directions and solve the feature extraction model with it. Compared with image domain sparsity constraint result, this time recovered image also has improved contrast and low resolution, but it doesn't contain additional ringing. (Fig.16, 17)

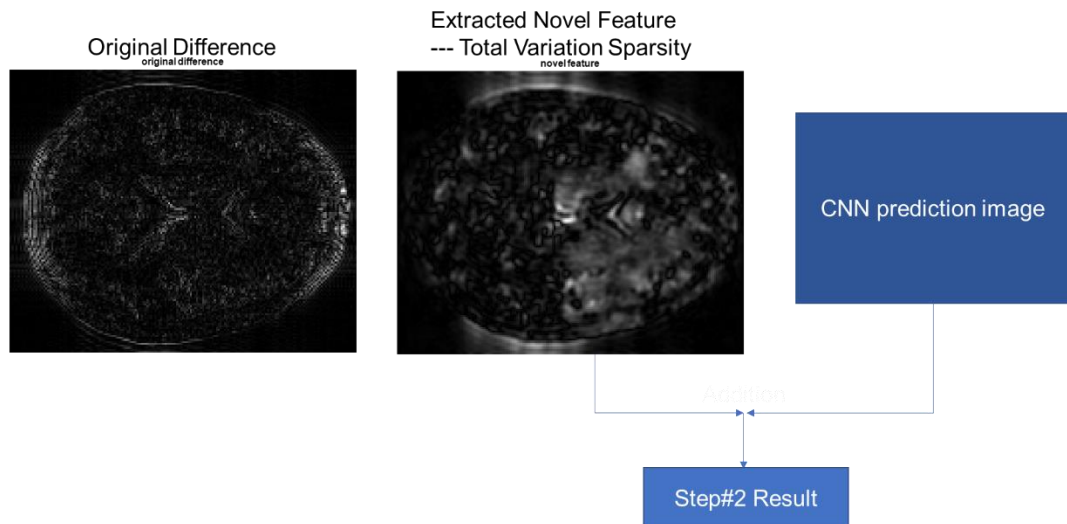


Fig.16 Extracted novel feature with TV sparsity

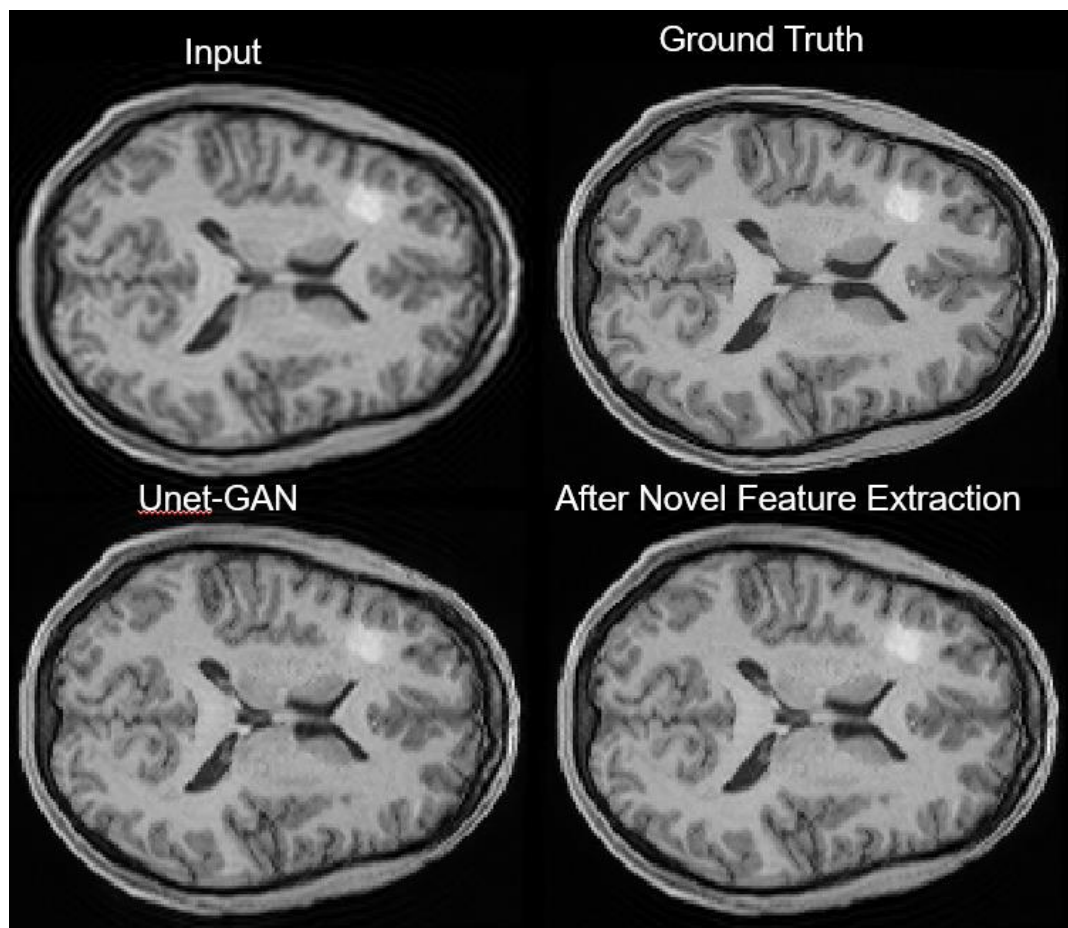


Fig.17 Novel Feature Extraction result with TV sparsity constraint

4.2.3 The improvement and validation

To solve the low-resolution problem, our future plan is

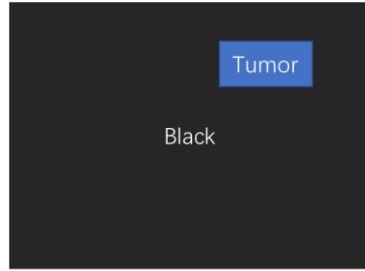
- 1. To locate novel feature

- 2. To get a realistic image of brain without feature in CNN

Our validation about these two goals is as follows.

We aim to extract novel features from inner k space difference.

Take the simplest case, we know the location and the tumor's low-resolution image. (Fig.18) We put the image into model (5) and limit its variation to known location. Then we get our result in Fig.19. It looks similar to ground truth with a little variation.



Low Resolution Image of Tumor

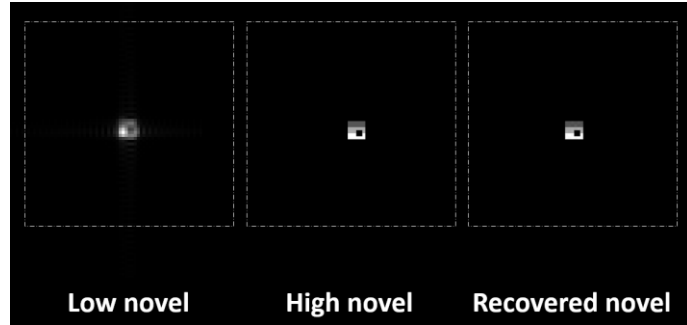
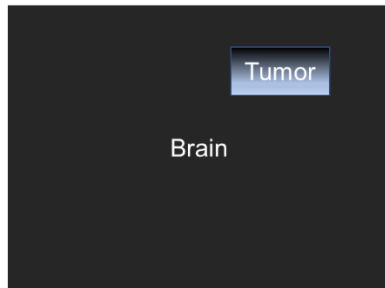


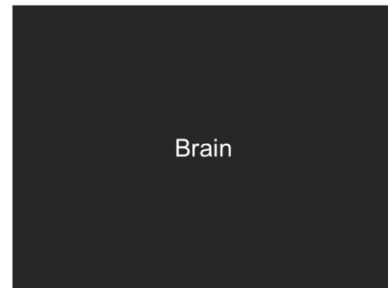
Fig.19 Stage 1 Result

Fig.18 Validation Stage 1

Then things become a little more complex. This time we know tumor's location, but we only have low resolution image of brain with tumor and low-resolution image without tumor. That means, surrounding pixels are warped into the tumor area. We also get a clear result in Fig.21.



Low Resolution Image of Brain with Tumor



Low Resolution Image without Tumor

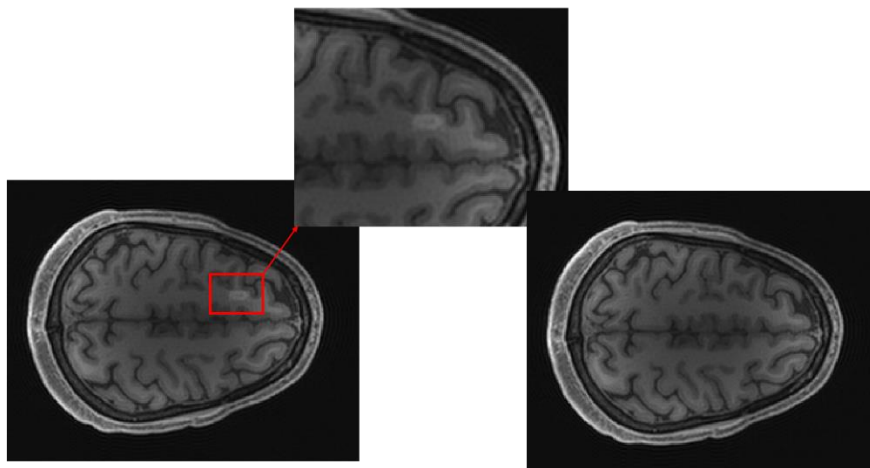


Fig.20 Validation Stage 2

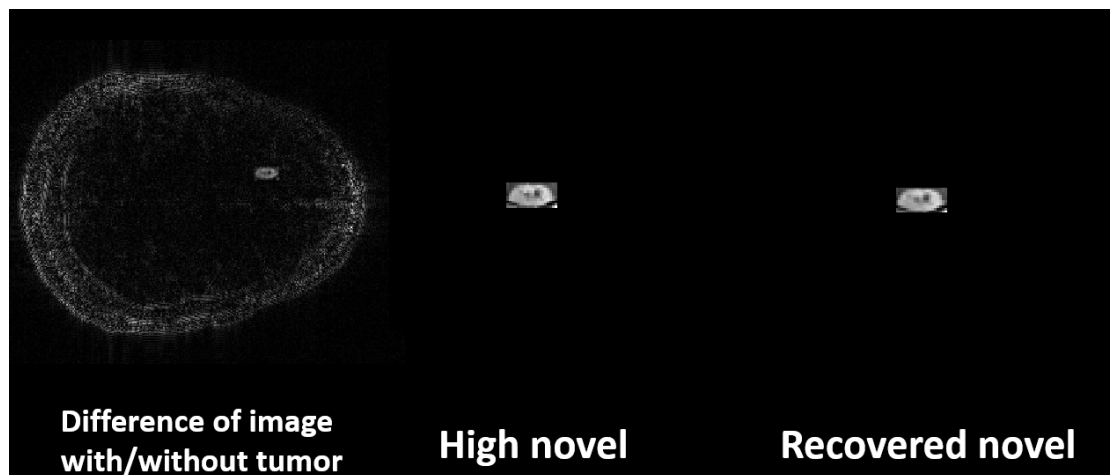
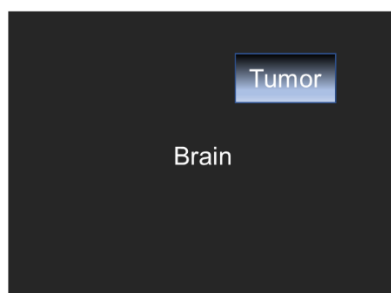
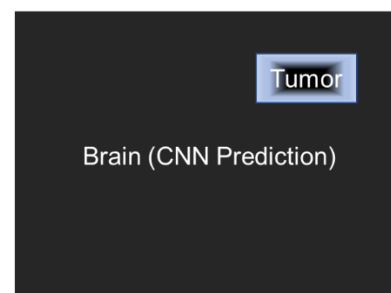


Fig.21 Stage 2 Result

And then, if we don't have low-resolution image without tumor, instead, we have CNN predicted image, things become rather tricky. If our CNN is ideal, its prediction should be very similar to brain image without tumor. In practical situation, there still exists some feature in image. Therefore, this time our result is not so good as before.



Low Resolution Image of Brain with Tumor



Low Resolution Image of CNN Prediction

Fig.22 Validation Stage 3

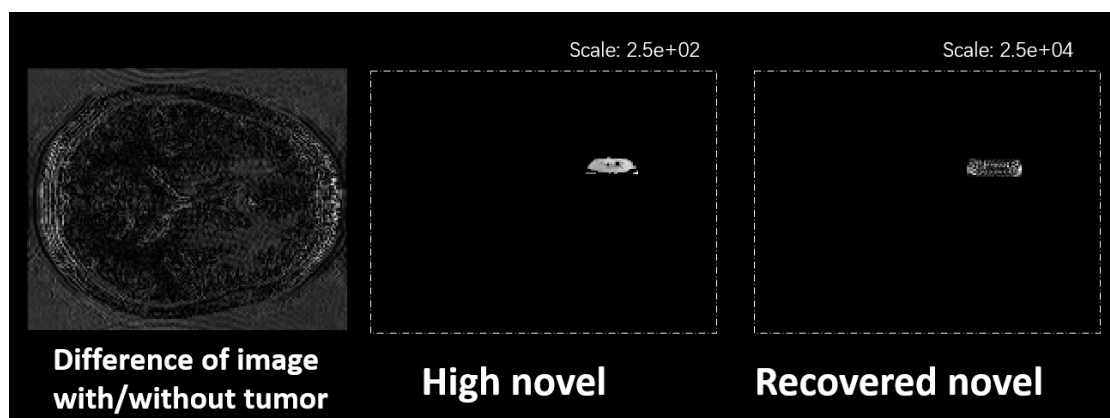


Fig.23 Stage 3 Result

And if we restrict the conditions, we don't know the location, it becomes our practical situation. We

don't know the location, so utilize sparsity instead; we don't have real brain image without tumor, so use low-resolution CNN prediction as an alternative. That's the reason why we cannot get a good result as stage1 or 2 does. And that's what we'll make effort to in the future.

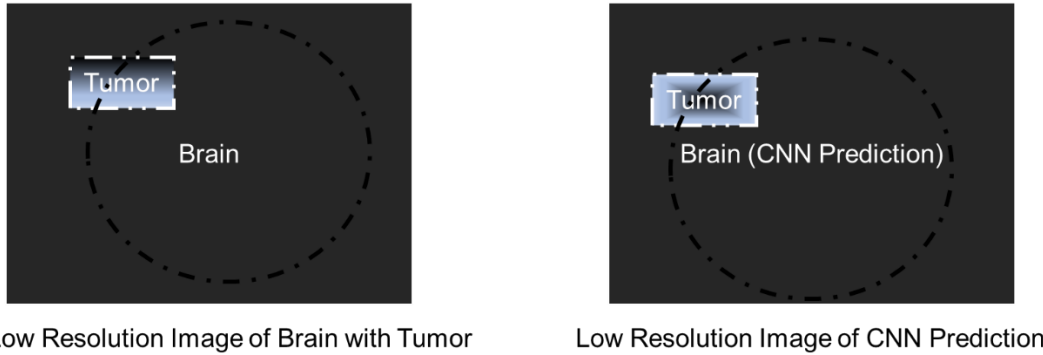


Fig.23 Practical Situation

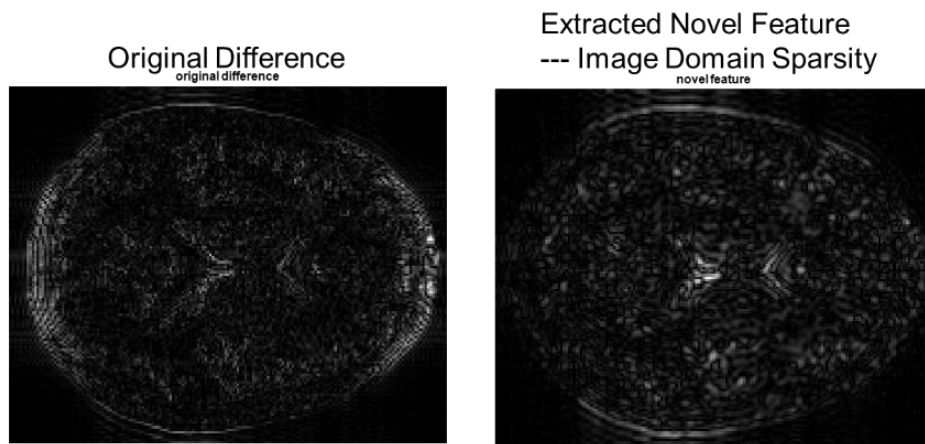


Fig.24 Current Result

4.3 Step#3 Optimization based on Data Consistency

Our conclusion in this step is:

1. If we utilize image domain sparsity in Step#2, it's necessary and helpful to implement Step#3, while if we utilize TV sparsity, it's not necessary.
2. Traditional way to do data consistency does little work in our case, while GSM performs well in novel feature that done in image domain but decreases the quality of image when using total variance because the latter way already performs well in novel feature.

4.3.1 Traditional Data Consistency Model Result

We can observe little improvement in assessments before and after optimization. (Fig.25,26)

Results of traditional way (image domain)

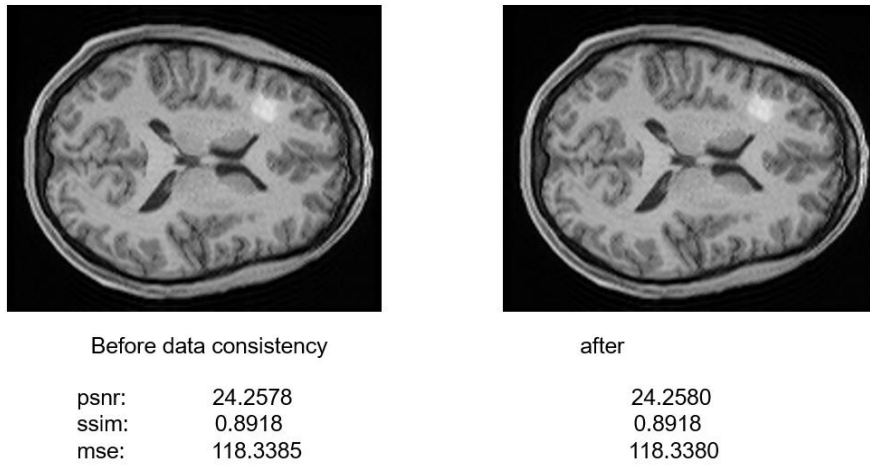


Fig.25 Traditional Model Result with extracted novel feature by image domain sparsity

Results of traditional way (total variance)

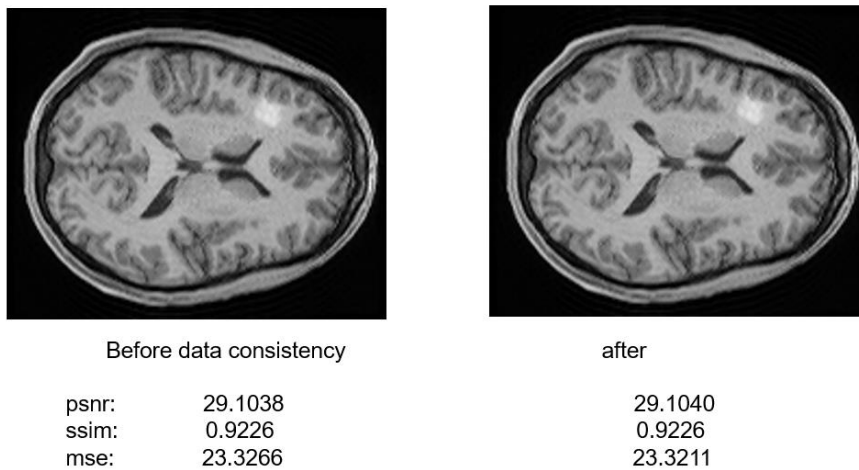


Fig.26 Traditional Model Result with extracted novel feature by TV sparsity

4.3.2 Generalized Series Model Result

We can observe obvious improvement in assessments before and after optimization by GS model with image domain sparsity, but detriment for that with TV sparsity. (Fig.27,28)

Results of GSM (image domain)

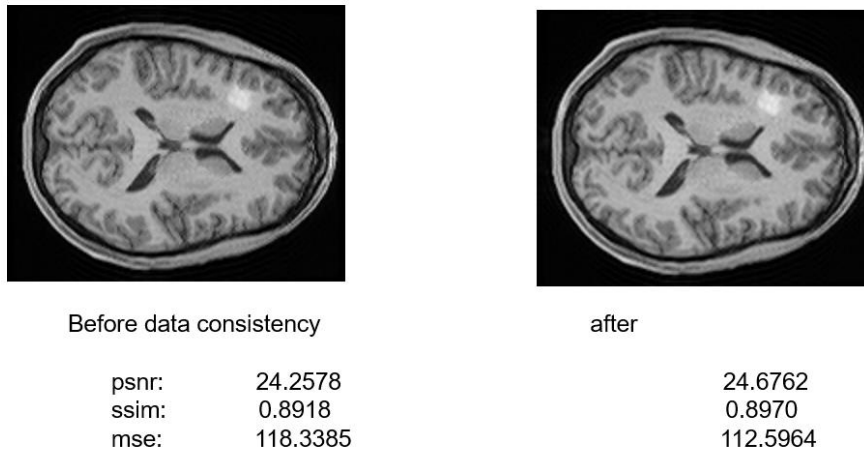


Fig.27 GS Model Result with extracted novel feature by image domain sparsity

Results of GSM (total variance)

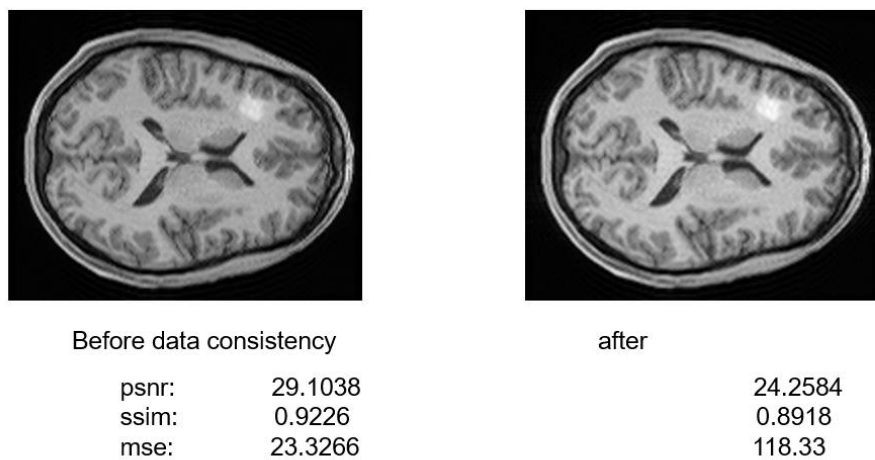


Fig.28 GS Model Result with extracted novel feature by TV sparsity

Conclusion

- We build a new strategy for super-resolution in MRI and achieve 4 times high resolution, with recovered novel feature and no artifact.
- For Deep Learning part, we utilize image normalization and no patch strategy to help neural network learn more efficiently and accurately.

- For Novel Feature Extraction part, we build model based on image sparsity and TV sparsity and recover novel features contrast.
- For Optimization based on Data Consistency part, we build two models based on Fourier Series and Generalized Series, and improve the image quality.
- Our future plan is to optimize neural network structure and find more suitable sparsity constraint and better data consistency.

- [1] E. V. Reeth, I. W. K. Tham, C. H. Tan, and C. L. Poh, "Super - resolution in magnetic resonance imaging: A review," *Concepts in Magnetic Resonance Part A*, vol. 40A, pp. 306-325, 2012.
- [2] S. Geethanath, R. Reddy, A. S. Konar, S. Imam, R. Sundaresan, and R. Venkatesan, "Compressed sensing MRI: a review," *Critical Reviews™ in Biomedical Engineering*, vol. 41, 2013.
- [3] A. Loktyushin, H. Nickisch, R. Pohmann, and B. Schölkopf, "Blind retrospective motion correction of MR images," *Magnetic resonance in medicine*, vol. 70, pp. 1608-1618, 2013.
- [4] J. Maclaren, M. Herbst, O. Speck, and M. Zaitsev, "Prospective motion correction in brain imaging: a review," *Magnetic resonance in medicine*, vol. 69, pp. 621-636, 2013.
- [5] M. T. McCann, K. H. Jin, and M. Unser, "Convolutional Neural Networks for Inverse Problems in Imaging: A Review," *IEEE Signal Processing Magazine*, vol. 34, pp. 85-95, 2017.
- [6] Y. Chen, Y. Xie, Z. Zhou, F. Shi, A. G. Christodoulou, and D. Li, "Brain MRI super resolution using 3D deep densely connected neural networks," in *2018 IEEE 15th International Symposium on Biomedical Imaging (ISBI 2018)*, 2018, pp. 739-742.
- [7] C. Ledig, L. Theis, F. Huszár, J. Caballero, A. Cunningham, A. Acosta, *et al.*, "Photo-Realistic Single Image Super-Resolution Using a Generative Adversarial Network," in *CVPR*, 2017, p. 4.
- [8] S. Mohamed, "A statistical view of deep learning," ed: Technical Report (blog. shakirm. com), 2015.
- [9] J. Ashburner, "A fast diffeomorphic image registration algorithm," *NeuroImage*, vol. 38, pp. 95-113, 2007/10/15/ 2007.
- [10] H. Bin, C. Weihai, W. Xingming, and L. Chun-Liang, "High-Quality Face Image SR Using Conditional Generative Adversarial Networks," *arXiv preprint arXiv:1707.00737*, 2017.
- [11] Z.-P. Liang and P. C. Lauterbur, "A generalized series approach to MR spectroscopic imaging," *IEEE Transactions on Medical imaging*, vol. 10, pp. 132-137, 1991.