

Chuyun SHEN, Wenhao LI, Qisen XU, Bin HU, Bo JIN, Haibin CAI, Fengping ZHU, Yuxin LI, Xiangfeng WANG, 2023. Interactive medical image segmentation with self-adaptive confidence calibration. *Frontiers of Information Technology & Electronic Engineering*, 24(9):1332-1348. <https://doi.org/10.1631/FITEE.2200299>

Interactive medical image segmentation with self-adaptive confidence calibration

Key words: Medical image segmentation; Interactive segmentation; Multi-agent reinforcement learning; Confidence learning; Semi-supervised learning

Corresponding author: Xiangfeng WANG

E-mail: xfwang@cs.ecnu.edu.cn

 ORCID: <https://orcid.org/0000-0003-1802-4425>

Motivation

Interactive medical image segmentation based on human-in-the-loop machine learning is a novel paradigm that draws on human expert knowledge to assist medical image segmentation.

However, existing methods often fall into what we call interactive misunderstanding, the essence of which is the dilemma in trading off short- and long-term interaction information.

To better use the interaction information at various timescales, we propose an interactive segmentation framework, called interactive MEdical image segmentation with self-adaptive Confidence CALibration (MECCA), which combines action-based confidence learning and multi-agent reinforcement learning.

Contributions

- A novel *confidence network* is learned by predicting the alignment level of the action with short-term interactive information.
- A confidence-based *reward shaping* mechanism is proposed to explicitly incorporate the confidence into policy gradient calculation, thus avoiding the model's *interactive misunderstanding*.
- The confidence network can generate high-quality labels for unlabeled data, thus reducing interaction *intensity*.
- One can generate advice interactive regions based on the confidence network to reduce interaction *difficulty*.

Framework

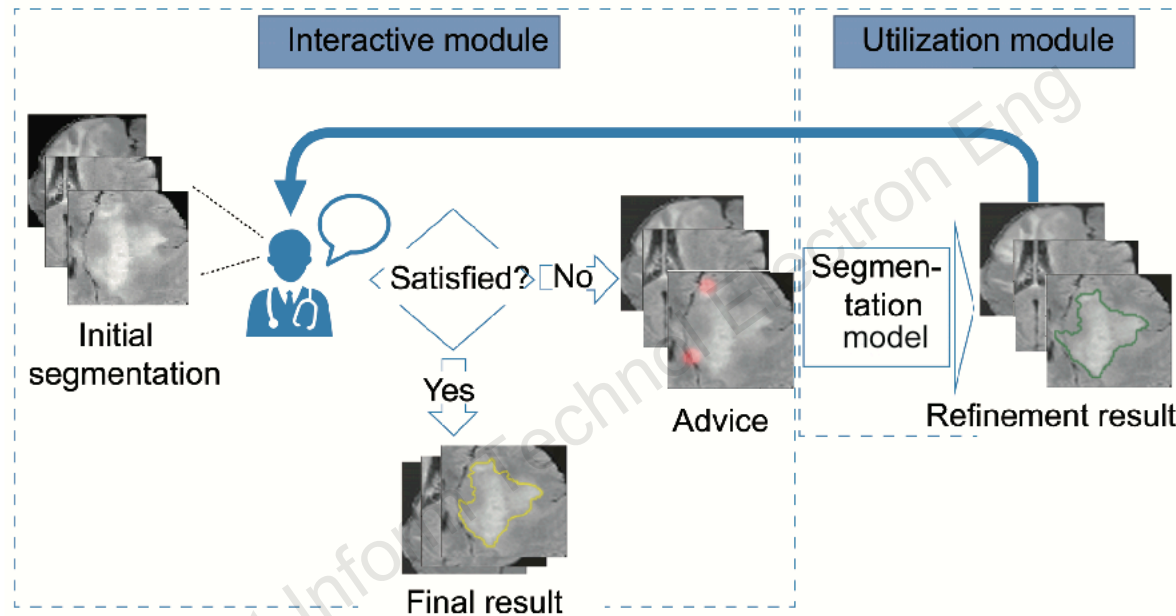


Fig. 1 Interactive segmentation process. In the interactive module, the expert observes the current (or initial) segmentation and provides correction information (the red hint); in the utilization module, new segmentation is refined based on the correction information. References to color refer to the online version of this figure

Framework

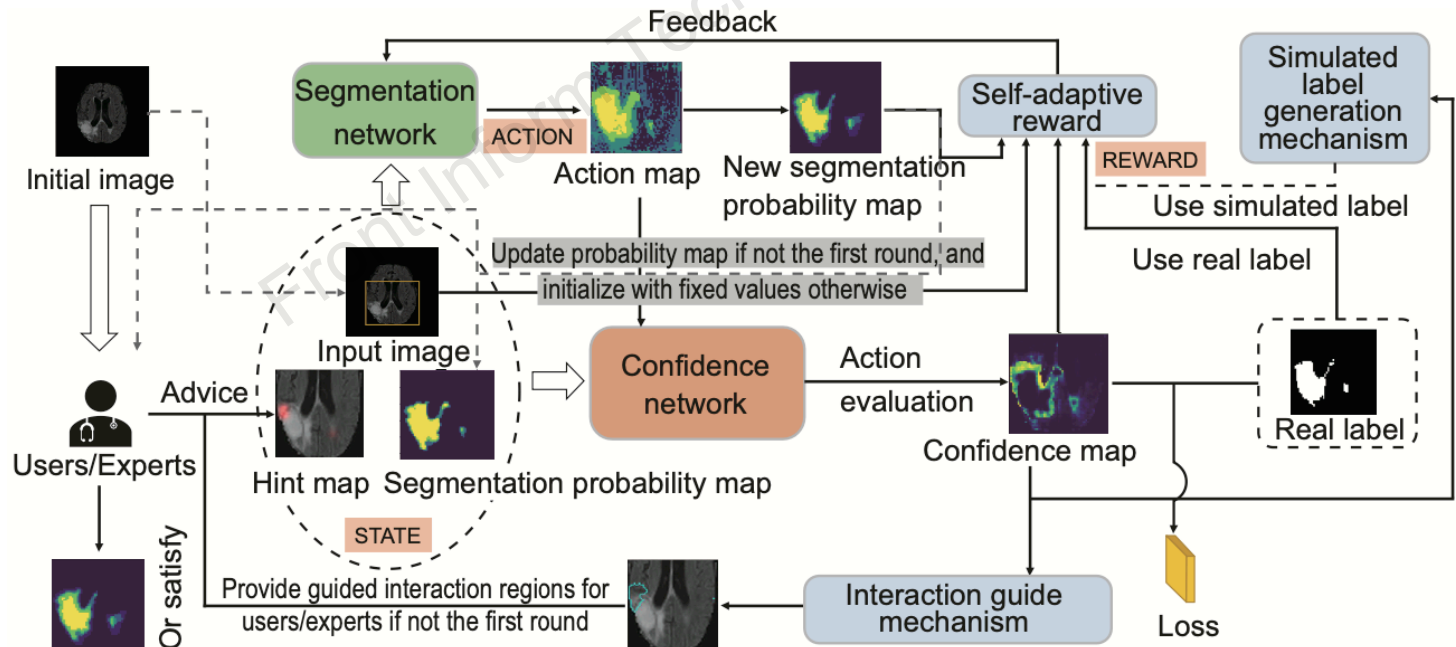
$x = (x_1, \dots, x_N)$ denotes the input image and x_i denotes the i^{th} voxel. Every voxel x_i is treated as an agent with its own refinement policy $\pi_i(a_i^{(t)}, s_i^{(t)})$. The state $s_i^{(t)}$ for agent x_i is concatenated by its voxel value b_i , its current segmentation probability $p_i^{(t)}$, and the value $h_i^{(t)}$ on the hint map. The action $a_i^{(t)}$ for agent x_i is sampled from its policy and used to adjust its previous segmentation probability:

$$\begin{aligned} a_i^{(t)} &\sim \pi_{\theta} \left(a_i^{(t)} \mid s_i^{(t)} \right), \\ p_i^{(t+1)} &= \text{clip} \left(p_i^{(t)} + a_i^{(t)}, 0, 1 \right). \end{aligned}$$

After taking the action, the agent will receive a reward $r_i^{(t)}$ according to the segmentation result.

Method

The segmentation module outputs actions to change the segmentation probability of each voxel (agent) at each interaction step. Meanwhile, the confidence network estimating the confidence of the actions will generate the self-adaptive reward and simulated label. The confidence map can provide the advice regions of the next interaction step to experts. The confidence map can provide the advice regions of the next interaction step to experts.



Method

Action-based confidence network

C denotes the confidence network, w_C denotes the parameters of the confidence network, while L_{BCE} denotes the binary cross-entropy loss:

$$\begin{aligned} L_C(s^{(t)}, a^{(t)}; w_C) \\ = L_{\text{BCE}}(C(s^{(t)}, a^{(t)}), g^{(t)}) + L_{\text{BCE}}(C(s^{(t)}, -a^{(t)}), 1 - g^{(t)}), \end{aligned}$$

where

$$g^{(t)} = \begin{cases} 0 & \text{if } a^{(t)} \oplus y^{(t)} == 1 \\ 1 & \text{otherwise} \end{cases}$$

and $g^{(t)}$ means whether the direction of action is consistent with the label. For $a \oplus b$, the statement is true only if either $a > 0$ or $b > 0$, not both.

Method

Confidence-based reward shaping

We formulate action-aware learning as the self-adaptive reward function, $r^{(t)}$, to adapt the mechanism to multi-agent reinforcement learning (MARL) training:

$$r^{(t)} = \sum_{i=1}^N \alpha(2 - c_i)^\beta \text{gain}_i^{(t)},$$

where c_i is the value on the confidence map, and $\text{gain}_i^{(t)}$ denotes the relative gain of cross-entropy:

$$\text{gain}_i^{(t)} = \chi_i^{(t-1)} - \chi_i^{(t)}, \chi_i^{(t)} = -y_i \ln(p_i^{(t)}) - (1 - y_i) \ln(1 - p_i^{(t)}),$$

where $\chi_i^{(t)}$ denotes the cross-entropy between the current segmentation probability and the ground truth.

Method

Label generation and interaction guidance

The confidence network can generate high-quality labels for unlabeled data, thus reducing interaction intensity.

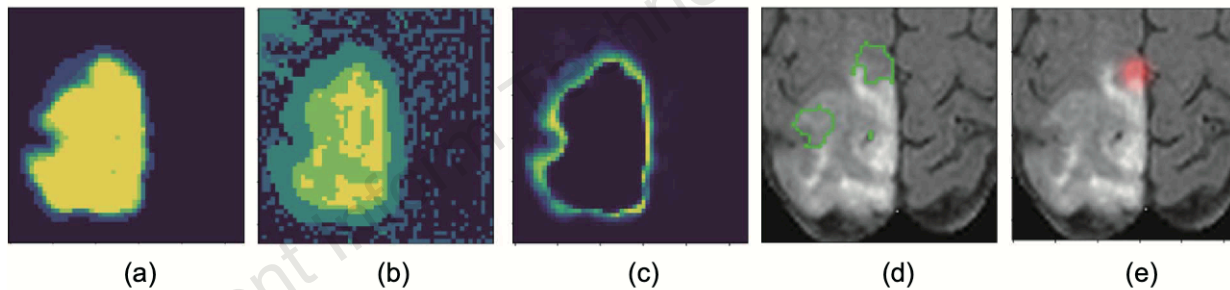


Fig. 5 Interaction guide mechanism: (a) probability map; (b) action map; (c) confidence map; (d) advice region; (e) interaction point. The areas surrounded by the green lines are advice regions provided to users, and the red point in (e) is the real hint information selected from advice regions. The brighter the color (closer to yellow), the larger the positive value; conversely, the darker the color (closer to black), the smaller the negative value. References to color refer to the online version of this figure

Major results

Table 1 Quantitative comparison of different methods

Method	Dice (%)				
	BraTS2020	BraTS2015	MM-WHS	Spleen	Liver
P-Net	83.67±8.35	84.00±12.01	81.40±1.48	88.08±2.25	35.89±2.61
U-Net	84.72±10.42	84.66±11.25	80.96±1.65	87.95±2.87	56.00±1.93
DeepIGeoS	88.54±.97	88.32±5.34	88.48±0.71	91.97±1.51	48.57±2.52
InterCNN	88.39±6.01	88.26±7.07	87.85±1.15	93.52±0.94	59.92±2.20
IteR-MRL	89.22±4.65	88.94±4.81	89.55±0.87	91.50±1.34	62.29±1.93
BS-IRIS	90.47±5.23	89.74±3.86	89.12±0.98	92.35±1.13	67.25±2.01
MECCA	91.02±5.86	90.29±5.07	90.39±5.89	94.96±1.44	71.46±1.41

Method	ASSD (pixel)				
	BraTS2020	BraTS2015	MM-WHS	Spleen	Liver
P-Net	5.78±4.01	5.10±3.72	3.28±0.45	4.25±2.07	6.05±4.18
U-Net	4.09±3.89	6.17±4.69	3.72±0.39	5.12±1.09	3.93±3.94
DeepIGeoS	2.11±1.30	2.28±1.24	1.53±0.18	0.93±0.46	1.80±0.69
InterCNN	2.01±1.09	1.81±2.09	0.80±0.15	0.54±0.83	1.04±0.56
IteR-MRL	2.07±0.91	1.41±0.22	0.90±0.11	0.67±0.21	0.87±0.59
BS-IRIS	1.82±0.33	1.61±0.42	1.19±0.16	0.54±0.19	0.76±0.24
MECCA	1.15±0.20	1.50±0.33	0.80±0.01	0.30±0.16	0.41±0.20

The best results are in bold

Major results

The performance of different interactive medical segmentation methods and methods with different weighting rewards

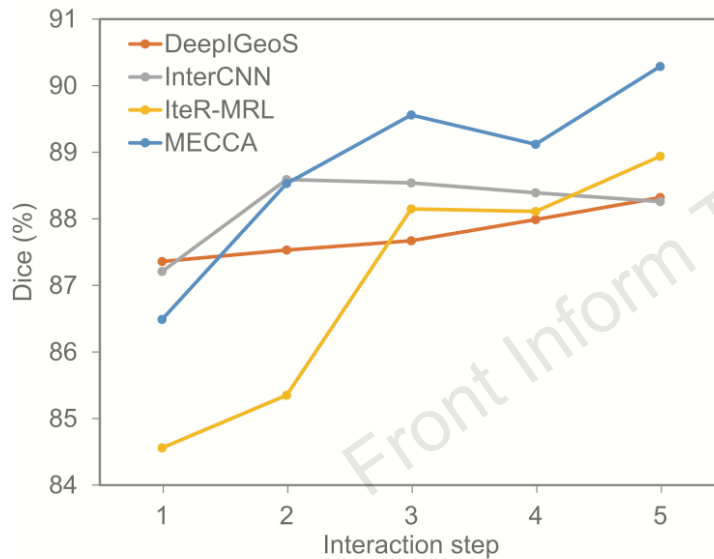


Fig. 8 The performance improvement of different interactive medical segmentation methods at different interaction steps. All these testing results were obtained on the BraTS2015 dataset

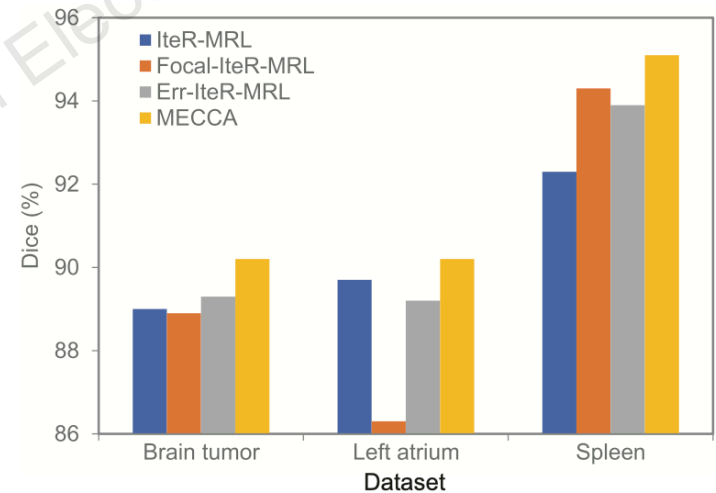


Fig. 9 Average Dice scores of methods with different weighting rewards. The reward function of IteR-MRL is not weighted; the reward function of Focal-IteR-MRL is weighted via Eq. (12), the reward function of Err-IteR-MRL is weighted in segmentation error regions, and the reward function of MECCA is weighted through action confidence

Conclusions

- We propose a method for learning the confidence of actions that continuously refine the segmentation result during the interaction process, so that hint information can be used more effectively.
- A self-adaptive reward is proposed for the segmentation module, which can prevent the misunderstanding during the interaction and help reduce the time cost of interaction by providing users with the advice regions with which they should interact next.
- MECCA is demonstrated to improve the efficiency of interactive medical image segmentation using fewer annotated samples.