

Multimodal Text and Image Classification



AutoML 연구실
석사과정 문아성
22.01.11

TO DO List

➤ Last week

- ☒ Bird Dataset: *"Fine-grained Image Classification via Combining Vision and Language"*
- ☒ 담낭용종 연구 실험 진행 (기존 서울대 데이터)
 - 기존 서울대 데이터로 adenoma, cancer, non-neoplastic 클래스 분류 실험 완료
- ☒ 졸업논문 – *"An Efficient Neural Network based on Early Compression of Sparse CT Slice Images"*

➤ This week

- ☒ Multimodal Text and Image Classification
- ☒ 담낭용종 연구 실험 진행 (기존 서울대 데이터)
 - Confusion matrix, ROC Curve 추가
- ☐ 졸업논문 – *"An Efficient Neural Network based on Early Compression of Sparse CT Slice Images"*
 - Introduction 내용 추가 (medical field efficiency의 중요성)
 - Related work 의학적, CNN 내용 추가 (CNN Algorithms)

Multimodal Text and Image Classification

- Deep Modular Co-Attention Networks for Visual Question Answering

Yu, Zhou, et al. "Deep modular co-attention networks for visual question answering."

Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. 2019.

- Self-Attention & Guided-Attention Units
- Modular Co-Attention (MCA)
- Modular Co-Attention Network

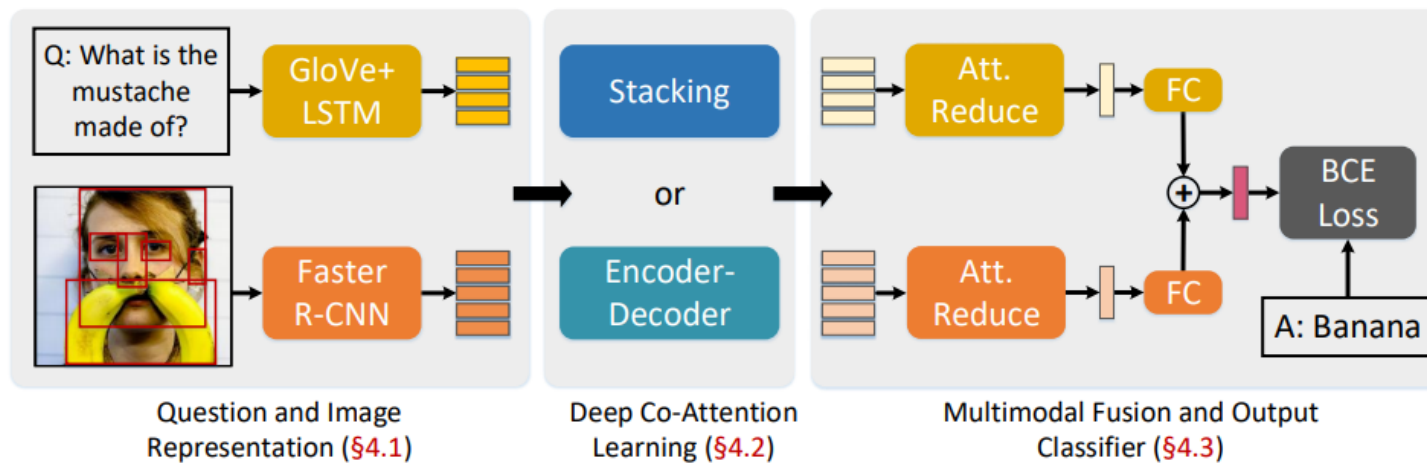


Figure 4: Overall flowchart of the deep Modular Co-Attention Networks (MCAN). In the Deep Co-attention Learning stage, we have two alternative strategies for deep co-attention learning, namely *stacking* and *encoder-decoder*.

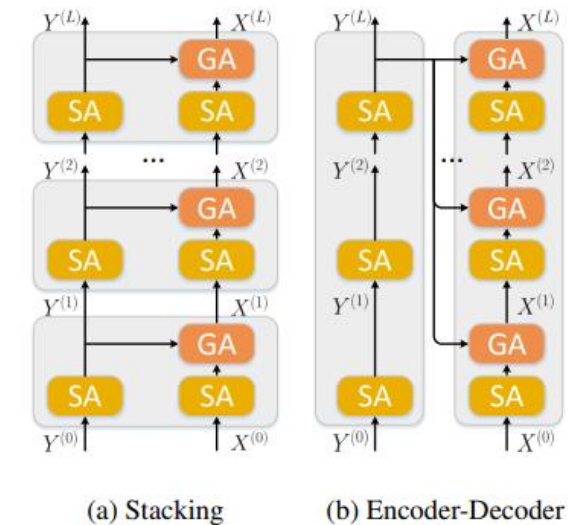


Figure 5: Two deep co-attention models based on a cascade of MCA layers (e.g., SA(Y)-SGA(X,Y)).

Multimodal Text and Image Classification

Deep Multimodal Neural Architecture Search

Yu, Zhou, et al. "Deep Multimodal Neural Architecture Search."

Proceedings of the 28th ACM International Conference on Multimedia. 2020.

Unified Encoder-Decoder Backbone

- Self-attention operation
- Guided-attention operation
- Feed-forward network operation
- Relation self-attention operation

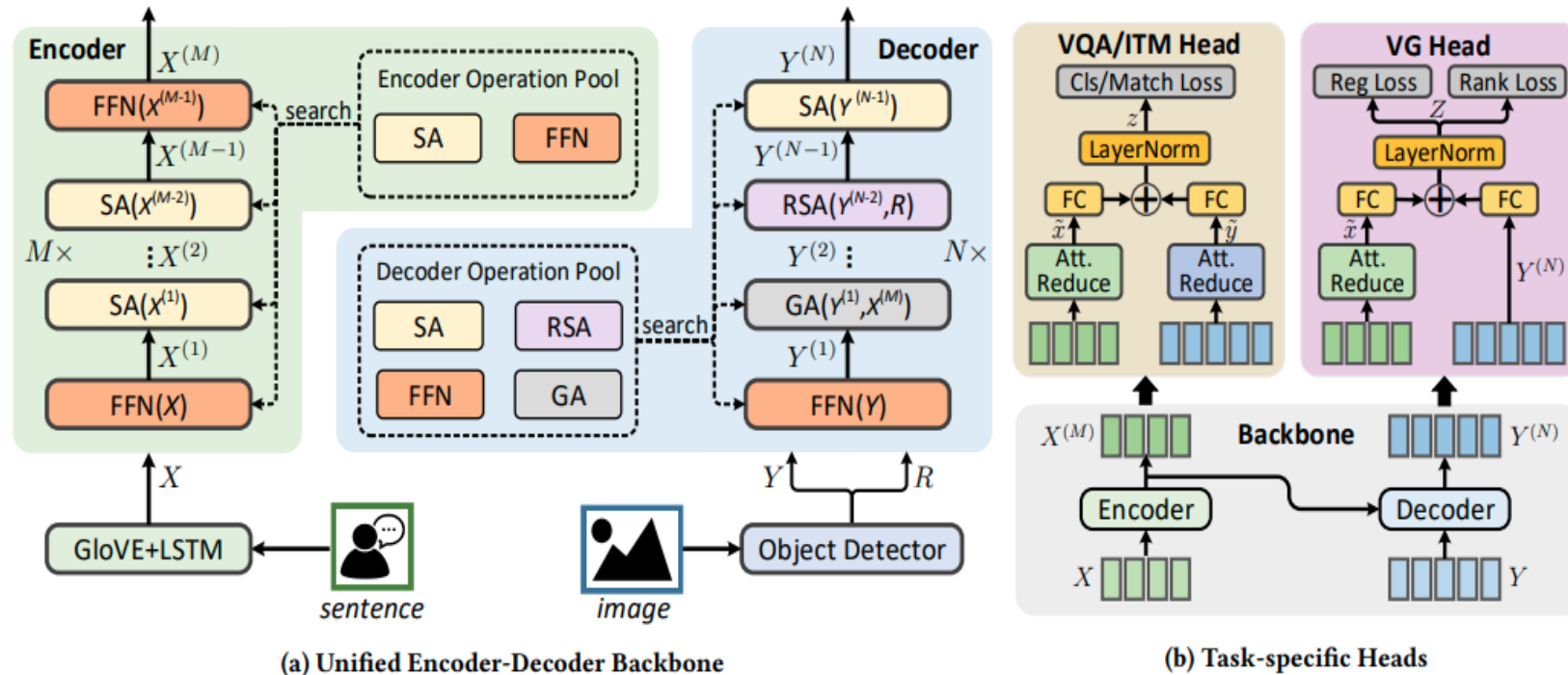


Figure 2: The flowchart of the MMnas framework, which consists of (a) unified encoder-decoder backbone and (b) task-specific heads on top the backbone for visual question answer (VQA), image-text matching (ITM), and visual grounding (VG).

Multimodal Text and Image Classification

■ Are These Birds Similar: Learning Branched Networks for Fine-grained Representations

Nawaz, Shah, et al. "Are these birds similar: Learning branched networks for fine-grained representations."
2019 International Conference on Image and Vision Computing New Zealand (IVCNZ). IEEE, 2019.

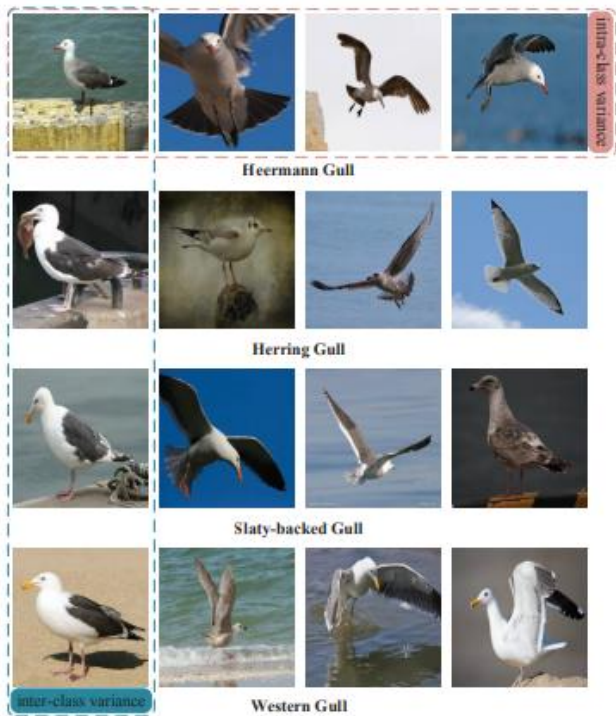


Fig. 1. Images extracted from CUB-200-2011. These examples indicate large intra-class variance and small inter-class variance, which makes the image classification task highly challenging.

Class	Vision	Natural Language Description
Green Jay		– This bird has a black pointed bill, with a black and green breast. – A colorful bird with shades of green, blue, and black. the beak is thick and pointed. – Bird has a bluish cheek patch and a black eyebrow with a bluish spot. ...
Common Yellowthroat		– This bird has a brown crown as well as a yellow throat. – This yellow breasted bird has a contrasting brown back and white flanks and thighs. – A small multi colored bird with a white breast and yellow back feathers. ...
Red Eyed Vireo		– This beautiful gold and gray colored bird had a sharp pointed beak and black tail. – This has a white underbelly with yellow and gray short feathers and a striped face. – A bird with a white belly and a dark brown eyebrow. ...

Fig. 3. Some random image-text pairs taken from the CUB-200-2011 dataset.

Multimodal Text and Image Classification

Model Architecture

- Vision Stream

*we exploited the visual information with the **NTS-NET** to extract features and perform the classification task.*

- Textual Stream

*we used only text descriptions and classified it with the **BERT model**.*

- Two Branch Network

we fused text and visual information with MLPs to perform multimodal classification.

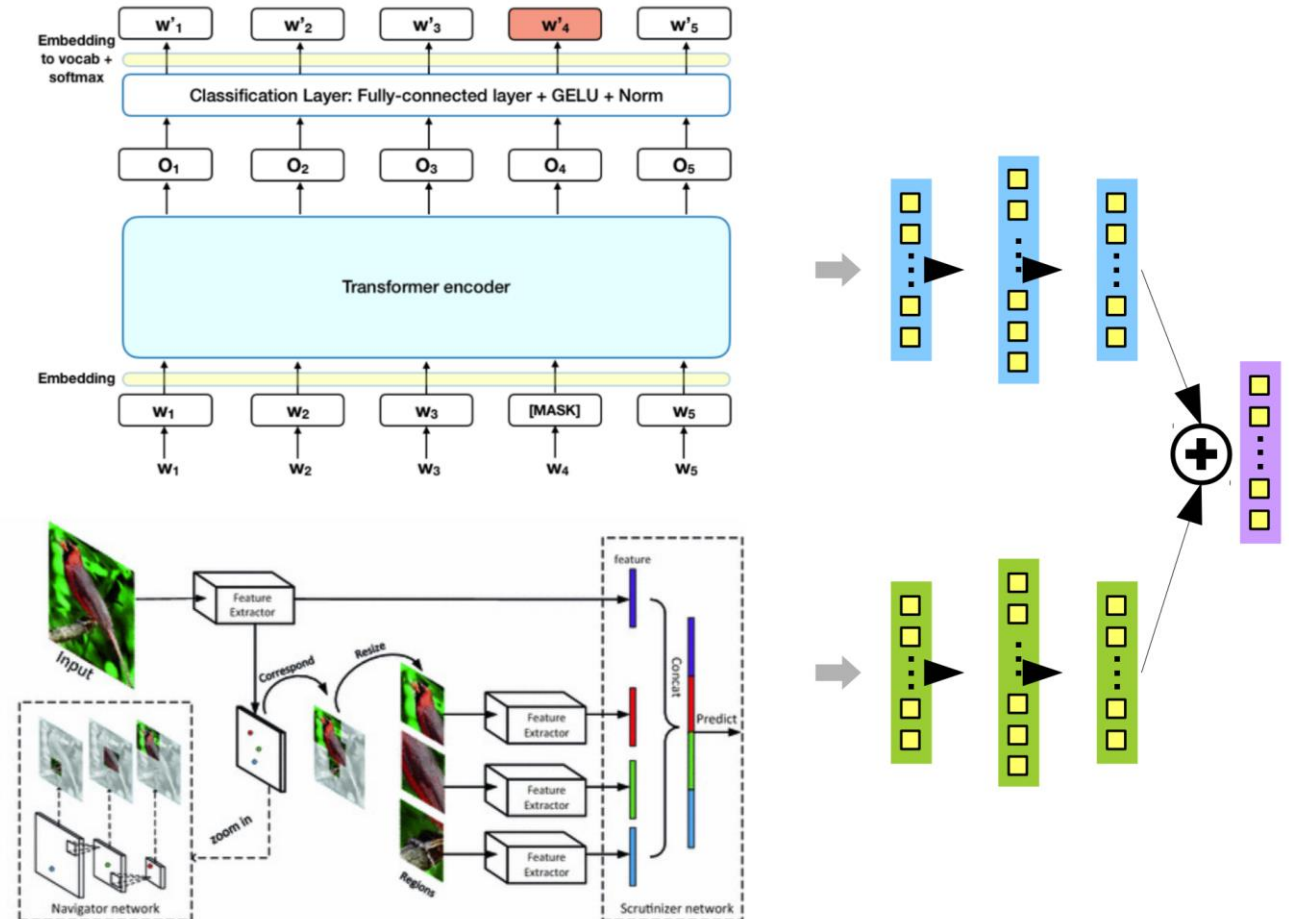


Fig. 2. The proposed approach extracts image and text representation from NTS-NET [9] and BERT [10] models. Finally, representations are fused after passing through a MLP.

Multimodal Text and Image Classification

- Image and Text fusion for UPMC Food-101 using BERT and CNNs

Gallo, Ignazio, et al. "Image and Text fusion for UPMC Food-101 using BERT and CNNs."

2020 35th International Conference on Image and Vision Computing New Zealand (IVCNZ). IEEE, 2020.



Examples of images in the UPMC Food-101 Dataset

10 Apple Pie and Cake Recipes

Holiday baking doesn't have to be bad for your waistline.

Pin it



1 OF 12

Dreamy desserts

Holiday baking doesn't have to be bad for your waistline.

Sweet or tart, apples are satisfying, thanks to the 4 grams of fiber in each serving. They're also packed with disease-fighting antioxidants.

These 10 apple desserts have lower fat, more fiber, and fewer calories than many traditional recipes, without sacrificing the flavors of the season.

Next: [Classic Apple Pie](#)

Text(recipe) data sample: Apple Pie

Multimodal Text and Image Classification

Model Architecture

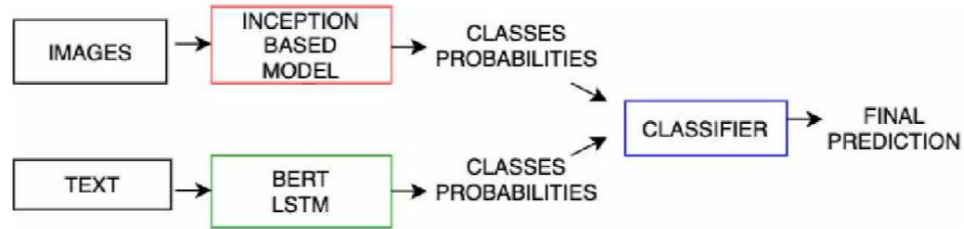


Fig. 1: Scheme of the implemented **late fusion** model with a stacking approach.

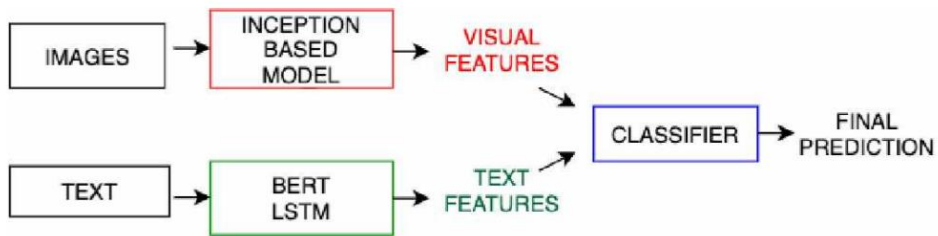


Fig. 2: Scheme of the implemented early fusion model with a stacking approach.

Input Image	Input Text
Inception V3	
Average Pooling 2D (8 x 8)	
Dropout - probability 40%	
Flatten	BERT Model
Dense - 128 units	LSTM - 128 units
Concatenate	
Dense - 256 units	
Dense <i>Softmax</i> - 101 units	

TABLE IV: Proposed model architecture.

기존 서울대 데이터 실험 세팅

➤ 데이터 구성

[기존 서울대 데이터]

- 177명의 환자 데이터 (전체 데이터)
- Cancer: 31명
- Adenoma: 69명
- Non-neoplastic : 77명

➤ Train & Test Set

[기존 서울대 데이터]

- 2,888장의 데이터셋
- Train Set : 141명(2,256장)
- Test Set : 36명(632장)

➤ Input Image

- Outline, Inline, Polyp, 담낭벽+Polyp 이미지

➤ 모델 구성

- 갑상선에서 사용한 모델을 담낭 용종 연구에 맞게 조정한 모델
- 데이터 불균형 문제를 해소하기 위한 Focal loss function 적용

➤ 반복 실험 30회

서울대 데이터 실험 결과

Table. The main classification metrics accuracy, AUC, sensitivity and specificity on a per-class basis.

	Accuracy	AUC	Sensitivity	Specificity
non-neoplastic	0.770	0.836	0.825	0.720
adenoma	0.715	0.805	0.519	0.833
cancer	0.821	0.784	0.586	0.897

AUC, area under the curve

서울대 데이터 실험 결과

Confusion matrix

<i>true labels</i>	non-neoplastic	0.35	0.06	0.02
	adenoma	0.11	0.18	0.06
	cancer	0.04	0.05	0.13
	<i>predicted labels</i>	non-neoplastic	adenoma	cancer

